

PHÂN TÍCH SỐNG CÒN

Giới thiệu

Trong Chương 15 chúng ta đã xem xét cách tính và phân tích tỉ lệ tử vong. Trong chương này chúng ta sẽ mô tả các phương pháp phân tích sống còn, chú trọng không chỉ vào sự kiện người đó có chết hay không mà còn vào thời gian sống **còn** (survival time) trước khi chết. Phương pháp phân tích sống còn có thể được áp dụng cho các biến cố như nghiên cứu thời gian bú sữa trong đó sự kiện kết thúc có thể là việc cai sữa.

Bảng sống

Mô thức sống còn của cộng đồng thường được mô tả trong bảng sống (life table). Bảng sống có 2 dạng. Dạng thứ nhất, bảng sống đoàn hệ (cohort life table). Trình bày thời gian sống còn của một nhóm cá nhân thực sự theo thời gian. Điểm khởi phát để đo thời gian sống còn có thể là lúc sinh hay có thể là biến cố khác. Thí dụ, một bảng sống đoàn hệ có thể dùng để trình bày tử vong của một nhóm nghề nghiệp tùy theo khoảng thời gian trong nghề, hay mô thức sống còn của bệnh nhân sau điều trị, thí dụ như sau khi điều trị, như xạ trị ung thư phế quản tế bào nhỏ (Bảng 16.1). Loại bảng sống thứ nhì, bảng sống hiện hành (current life table) thường được dùng trong mục đích bảo hiểm và ít phổ biến trong nghiên cứu y khoa. Bảng này trình bày thời gian sống còn kì vọng theo thời gian của một dân số giả thuyết được áp dụng tỉ suất tử vong đặc hiệu tuổi giới tính hiện nay.

Việc xây dựng bảng sống được mô tả nhờ xem xét bảng 16.1, trình bày thời gian sống còn bệnh nhân bị ung thư phế quản tế bào nhỏ, từng tháng một sau khi xạ trị. Bảng này dựa trên số liệu thu thập trong giai đoạn 5 năm. Số liệu được tóm tắt trong cột 1-4 của bảng sống, và các giá trị tính toán ở trong cột 5-8.

Cột 1 trình bày số tháng từ khi bắt đầu xạ trị. Cột 2 và 3 gồm số bệnh nhân còn sống ở đầu tháng và số người chết trong tháng. Thí dụ, 12 trong 240 bệnh nhân chết trong tháng thứ nhất sau khi điều trị, còn lại 228 người còn sống ở đầu tháng thứ hai.

Trong nghiên cứu loại này rất hiếm khi có thể theo dõi tất cả các bệnh nhân cho đến khi họ chết. Thứ nhất là do một số bệnh nhân có thể đi ra khỏi vùng nghiên cứu hay quyết định không còn tham gia. Điều này là một vấn đề quan trọng đối với một nghiên cứu theo dõi kéo dài. Thứ nhì bởi vì bệnh nhân thườn tham gia dần dần vào nghiên cứu, những người tham gia ở cuối thời gian nghiên cứu sẽ chỉ được theo dõi trong thời gian ngắn. Vì vậy trừ khi họ chết sớm, họ cũng thường không được theo dõi đủ. Do đó có những bệnh nhân sống tới một thời điểm nào đó nhưng tình trạng sống còn sau đó không rõ. Thời gian sống còn của họ gọi là được kiểm duyệt (censored). Những bệnh nhân đó được trình bày trong cột 4. Tổng số người nguy cơ bị chết trong tháng, sau khi đã được điều chỉnh cho sự thất lạc, được trình bày ở cột 5. Nó bằng với số bệnh nhân sống ở đầu tháng, trừ cho phân nửa số bị thất lạc, giả thiết rằng trung bình những sự thất lạc này xảy ra ở giữa tháng.

Bảng 16.1 Bảng sống trình bày kiểu hình sống còn của 240 bệnh nhân bị ung thư phế quản tế bào nhỏ được xạ trị

(1) Thời khoảng từ khi bắt đầu điều trị	(2) Số còn sống ở đầu thời khoảng	(3) Số chết trong thời khoảng	(4) Số không theo dõi được trong thời khoảng	(5) Số người nguy cơ = (2) - (4)/2	(6) nguy cơ chết trong thời khoảng	(7) Cơ hội sống còn trong thời khoảng	(8) Cơ hội sống còn lũy tích từ đầu trị liệu = (8,
1	240	12	0	240,0	0,0500	0,9500	0,9500
2	228	9	0	228,0	0,0395	0,9605	0,9125
3	219	17	1	218,5	0,0778	0,9222	0,8415
4	201	36	4	199,0	0,1809	0,8191	0,6893
5	161	6	2	160,0	0,0375	0,9625	0,6634
6	153	18	7	149,5	0,1204	0,8796	0,5835
7	128	13	5	125,5	0,1036	0,8986	0,5231
8	110	11	3	108,5	0,1014	0,8519	0,4700
9	96	14	3	94,5	0,1481	0,8986	0,4004
10	79	13	0	79,0	0,1646	0,8354	0,3345
11	66	15	4	64,0	0,2344	0,7656	0,2561
12	47	6	1	46,5	0,1290	0,8710	0,2231
13	40	6	0	40,0	0,1500	0,8500	0,1896
14	34	4	2	33,0	0,1212	0,8788	0,1666
15	28	5	0	28,0	0,1786	0,8214	0,1369
16	23	7	1	22,5	0,3111	0,6889	0,0943
17	15	12	0	15,0	0,8000	0,2000	0,0189
18	3	3	0	3,0	1,0000	0,0000	0,0000

Cột 6 trình bày nguy cơ tử vong trong tháng, được tính bằng số chết trong tháng chia cho số người nguy cơ. Cột 7 là cơ hội sống còn trong tháng, cột 8 là cơ hội sống còn chung (overall) hay tích lũy (cumulative), bằng cơ hội sống còn tính tới cuối tháng trước nhân với cơ hội sống còn trong tháng. Thí dụ, nguy cơ tử vong trong tháng thứ nhất là 12/240 hay 0,0500. Cơ hội sống còn do đó bằng 1 trừ 0,05 bằng 0,95. Trong tháng thứ hai 9 trong số 228 người còn lại bị chết tính ra nguy cơ tử vong là 0,0395 và cơ hội sống còn là 0,9605. Điều này có nghĩa là bệnh nhân đã sống tới tháng thứ nhất (với xác suất 0,9500) thì có cơ hội sống còn trong tháng thứ hai là 0,9605. Cơ hội sống còn chung hai tháng sau khi bắt đầu trị liệu là $0,9500 \times 0,9605 = 0,9125$.

Đường cong sống còn

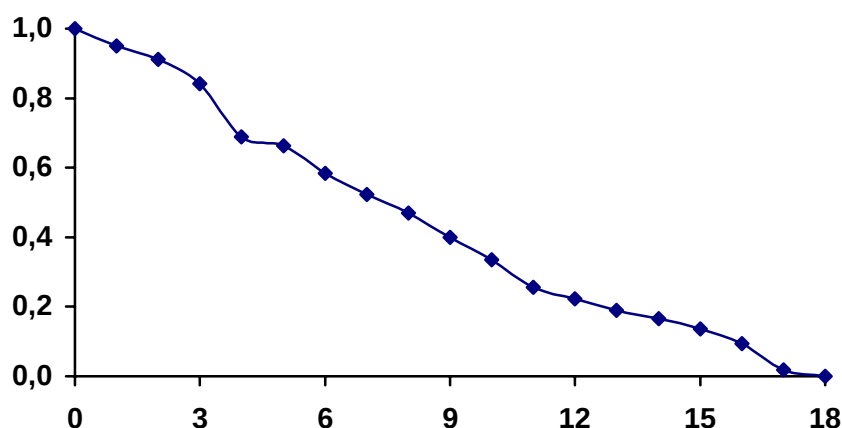
Trong nghiên cứu này tất cả những bệnh nhân chết sau 18 tháng. Đường cong sống còn vẽ từ các giá trị trong cột 8 được vẽ trong Hình 6.1.

Kì vọng sống

Thời gian sống còn trung bình hay kì vọng sống (life expectancy) cũng đáng được quan tâm. Có thể tính bằng cách dùng cột 1 và 8 của bảng sống. Trong mỗi thời khoảng, thời gian sống được nhân với cơ hội sống còn lũy tích. Tổng các giá trị này cộng thêm 1/2 là kì vọng sống. (cộng vào 1/2 là do tác động của việc phân nhóm bảng sống theo nguyên tháng và tương tự như hiệu chỉnh tính liên tục mà chúng ta đã gặp trong những chương trước).

$$K_{\text{vọng sống}} = \frac{1}{2} + \sum \left(\frac{\text{Số hàng trong mẫu thử khoảng} \times \text{cơ hội sống còn lũy tích}}{\text{mẫu thử khoảng}} \right)$$

Trong bảng 16.1 tất cả các khoảng là một tháng và do đó kì vọng sống là tổng các giá trị trong cột 8 cộng với 1/2 bằng 7,95 tháng.



Hình 16.1 Đường con sống còn cho bệnh nhân bị ung thư phế quản tế bào nhỏ được điều trị bằng xạ trị. Đường cong được vẽ từ bảng sống, bảng 16.1.

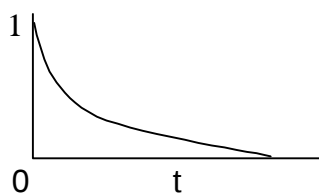
So sánh các bảng sống

Bước đầu tiên so sánh hai bảng sống bất kì là vẽ hai đường cong sống còn tương ứng và so sánh chúng bằng mắt. Ý nghĩa thống kê của sự khác biệt quan sát được đánh giá bằng kiểm định log rank. Đây là một ứng dụng đặc biệt của thao tác χ^2 Mantel Haenszel được tiến hành bằng cách xây dựng bảng 2×2 cho mỗi khoảng của bảng sống, để so sánh các tỉ lệ chết trong mỗi khoảng. Thí dụ, nếu bảng 16.1 cho thấy đường cong sống còn sau khi xạ trị ung thư phế quản tế bào nhỏ được so sánh với đường cong sống còn sau khi hóa trị, thủ tục này sẽ tạo ra 18 bảng 2×2 , mỗi bảng cho một tháng sau khi điều trị. Ứng dụng kiểm định χ^2 Mantel Haenszel cho những bảng này sẽ tổng kết sự khác biệt hàng tháng giữa số chết quan sát được sau khi xạ trị và số kì vọng nếu mô thức sống còn sau khi xạ trị giống như sau khi hóa trị. Kiểm định log rank có thể mở rộng để điều chỉnh cơ cấu khác nhau của nhóm điều trị, như tỉ số giới tính khác nhau hay phân phối tuổi khác nhau. thí dụ, giới tính sẽ được tính đến trong thí dụ này bằng cách xây dựng bảng 2×2 cho nam và nữ cho mỗi khoảng của bảng sống, và áp dụng kiểm định χ^2 Mantel Haenszel cho 36 (2×18) bảng. Để biết thêm chi tiết, xem bài báo của Peto et al (1976, 1977), hướng dẫn thực hành cho việc phân tích số liệu sống còn.

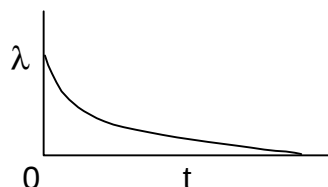
Một phương pháp khác, tinh vi hơn, để so sánh đường cong sống còn là dùng mô **hình nguy cơ tỉ lệ của Cox (Cox's proportional model)**, một **phương pháp** giống như hồi quy bội cho số trung bình và hồi quy logistic cho tỉ lệ. Nó được gọi là mô hình nguy cơ bởi vì nó giả định rằng tỉ lệ của nguy cơ chết trong hai nhóm hằng định theo thời gian, và tỉ lệ này giống nhau cho mỗi nhóm số liệu nhỏ, như các nhóm tuổi- giới tính khác nhau. Xem chi tiết ở Cox và Oakes (1984).



(a) Nguy cơ tử vong = hằng số, λ



(b) Đường cong sống còn: tỉ lệ sống còn = $e^{-\lambda t}$



(c) Phân phối của thời gian sống còn = $\lambda e^{-\lambda t}$ thời gian sống trung bình = $1/\lambda$

Hình 16.2 Mô hình mũ: nguy cơ chết theo thời gian (t) hằng định

Mô thức sống còn

Xem cột 6 của bảng 16.1 cho thấy rõ nguy cơ tử vong của ung thư tế bào nhỏ gia tăng với thời gian sau khi xạ trị. Người ta đã quan tâm nhiều đến việc tìm kiếm các hàm số toán học mô tả sự liên quan giữa nguy cơ chết (hay còn gọi là tỉ suất **nguy hại - hazard rate**) theo thời gian. **Mô hình đơn giản nhất và nổi tiếng nhất là mô hình mũ (exponential)**, trong đó nguy cơ chết hằng định theo thời gian. Có thể chứng minh rằng nếu một dân số bị một nguy cơ chết hằng định, không những đường còn sẽ có dạng mũ như hình 16.2 mà phân phối của thời gian sống còn cũng có dạng hình mũ với kì vọng sống bằng nghịch đảo của nguy cơ chết. Mô hình mũ thích hợp khi xem xét đời sống của các tạo vật có đời sống ngắn như muỗi, nguy cơ chết chính không phải vì tuổi tác mà là bởi vì tai nạn trong môi trường. Và nó đã được dùng trong việc xây dựng mô hình toán học truyền bệnh sốt rét. Các mô hình khác có các hàm số gamma, Weibull, Rayleigh và Gompertz để mô tả nguy cơ chết thay đổi theo thời gian. Mặc dù chúng hữu ích trong việc mô hình hóa bệnh tật và sống còn, những mô hình này không hữu ích trong việc phân tích số liệu sống còn bằng bảng sống và các phương pháp liên quan. Để biết chi tiết về các mô hình khác xem Cox và Oakes (1984).

PHÂN PHỐI POISSON

Giới thiệu

Chúng ta đã gặp phân phối bình thường cho trung bình và phân phối nhị thức cho tỉ lệ. Trong chương này ta gặp phân phối Poisson, được đặt tên theo tên một nhà toán học Pháp, rất thích hợp cho việc mô tả số lần xuất hiện biến cố theo thời gian, với điều kiện những biến cố này độc lập với nhau và ngẫu nhiên. Một thí dụ là số lần bức xạ tia phóng xạ được phát hiện bởi máy đếm lấp lánh (scintillation counter) trong 5 phút (xem Chương 17.2). Sau khi trình bày tổng quát phân phối Poisson, chúng ta sẽ xét các ứng dụng đặc biệt để phân tích tỉ suất, kể cả tỉ suất mới mắc bệnh.

Phân phối Poisson cũng thích hợp cho các hạt tìm thấy trong một đơn vị không gian, như là số các kí sinh trùng sốt rét trong kính hiển vi của một lam máu, với điều kiện các hạt phân phối ngẫu nhiên và độc lập trên toàn bộ không gian. Muốn đạt được phân phối Poisson cần thỏa hai thuộc tính ngẫu nhiên và độc lập. Thí dụ số trứng *Schistosoma mansoni* trong mẫu phân trong phải là phân phối Poisson bởi vì trứng có khuynh hướng dích chùm chứ không phải phân phối độc lập.

Ở Chương 18 chúng ta sẽ mô tả cách đánh giá một nhóm số liệu có tuân theo phân phối Poisson hay không và thí dụ 18.2 trình bày số liệu không tuân theo thuộc tính ngẫu nhiên. Hơn nữa, đã có những kĩ thuật để xác định các bệnh có tập trung trong không gian và thời gian hay không (xem tổng quan của Smith, 1982), như là khả năng tập trung các trường hợp bệnh bạch huyết ở trẻ quan các nhà máy năng lượng hạt nhân. Tập trung như vậy vi phạm phân phối Poisson.

Định nghĩa

Phân phối Poisson mô tả phân phối lấy mẫu của số lần xuất hiện, r , của biến cố trong một khoảng thời gian (hay vùng không gian). Nó phụ thuộc chỉ vào một tham số đó là số lần xuất hiện trung bình, m , trong khoảng thời gian bằng nhau (hay trong vùng không gian bằng nhau).

$$\text{Xác suất}(r \text{ lần xuất hiện}) = \frac{e^{-\mu} \mu^r}{r!}$$

Hầu hết các máy tính cầm tay đều có phím cho 3 hàm số cơ bản của công thức này. Lưu ý rằng, theo định nghĩa, cả $0!$ và μ^0 bằng 1. Do đó xác suất không xuất hiện biến cố là $e^{-\mu}$.

$$s.e. = \sqrt{\mu}$$

Sai số chuẩn cho số lần xuất hiện bằng căn bậc hai của trung bình, được ước lượng bằng căn bậc hai của số lần xuất hiện quan sát, (r) .

Thí dụ 17.1

Nhà chức trách y tế tại một quận có kế hoạch đóng cửa một trong hai phòng sinh đánh giá nhu cầu gia tăng đặt ra cho phòng sanh còn lại. Một yếu tố cần xem xét là nguy cơ vào một ngày nào đó nhu cầu nhập viện vượt quá khả năng của phòng. Hiện nay phòng sinh lớn có trung bình 4,2 lần nhập viện trong ngày và có khả năng giải quyết 10 lần nhập viện mỗi ngày. Sau khi đóng cửa phòng sinh nhỏ số lần nhập viện trung bình sẽ lên khoảng 6,1 lần mỗi ngày. Phân phối Poisson được dùng để ước lượng tỉ lệ số ngày mà khả năng của phòng bị quá tải. Thí dụ ta cần xem xác suất nhập viện hơn hoặc bằng 11 lần mỗi ngày. Điều này có thể tính được bằng cách tính xác suất có 0, 1, 2,... tới 10 lần nhập viện và trừ tổng số cho 1, như tính trong bảng 17.1. Thí dụ:

$$\text{Xác suất}(3 \text{ lần nhập viện}) = \frac{e^{-6,1} 6,1^3}{3!} = 0,0848$$

Tính toán cho thấy xác suất nhập viện 11 lần hay hơn trong một ngày là 0,0470. Do đó khả năng của phòng có thể bị quá tải 4,7% lần hay khoảng 17 ngày trong năm.

Bảng 17.1 Xác suất số lần nhập viện trong ngày ở một phòng sinh, dựa trên phân phối Poisson với trung bình 6,1 lần mỗi ngày.

Số lần nhập viện	Xác suất
0	0,0022
1	0,0137
2	0,0417
3	0,0848
4	0,1294
5	0,1579
6	0,1605
7	0,1399
8	0,1066
9	0,0723
Tổng số (0-10)	0,9530
11+ (làm phép trừ)	0,0470

Hình dáng

Hình 17.1 vẽ hình dáng của phân phối Poisson cho các giá trị trung bình, μ . Phân phối này rất lệch khi trung bình nhỏ, khi đó có xác suất không xảy ra biến cố rất đáng kể. Nó đối xứng với trung bình lớn và có thể xấp xỉ bằng phân phối bình thường nếu $\mu \geq 10$.

Kết hợp số đếm

Thí dụ 17.2

Một bệnh phẩm được đánh dấu bằng một chất phóng xạ được đếm trong 5 phút trong một buồng đếm lấp lánh đếm được 2905. Sai số chuẩn của số này được ước tính bằng phân phối Poisson

$$\text{Số hạt đếm được} = 2905$$

$$\text{s.e.} = \sqrt{2905} = 53,9$$

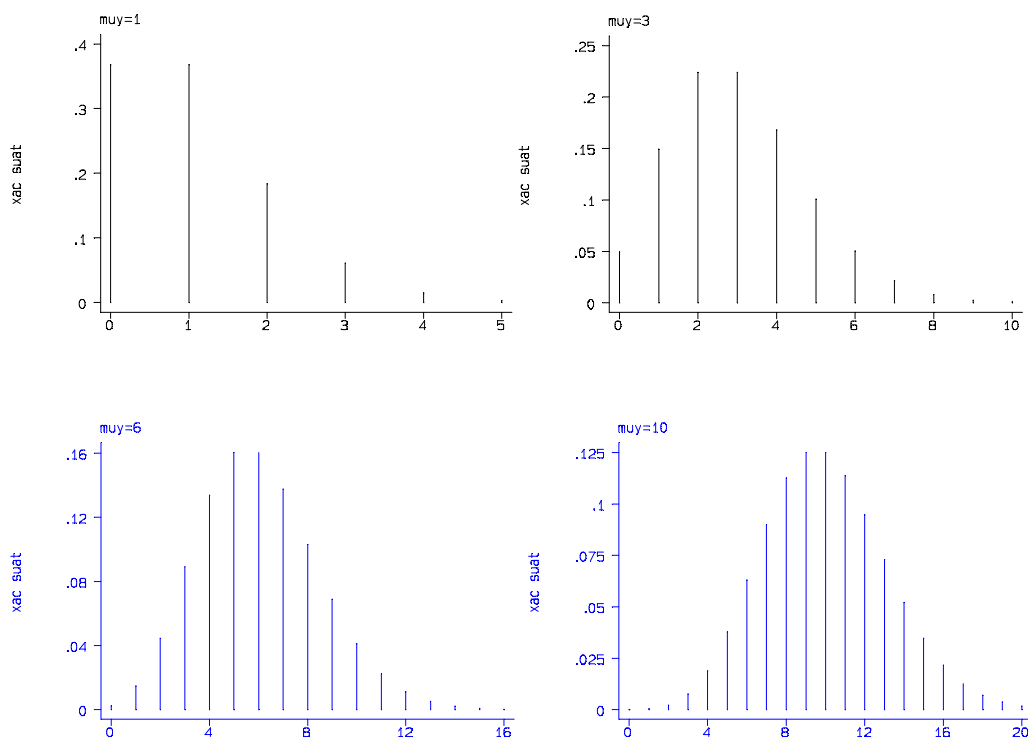
Người ta thường trình bày kết quả theo số đếm/phút

$$\text{Số hạt đếm được trong phút} = 2905/5 = 581,0$$

$$\text{s.e.} = 53,9/5 = 10,8$$

Khoảng tin cậy trung bình số hạt đếm được trong một phút được tính bằng cách dùng xấp xỉ bình thường:

$$581 \pm 1,96 \times 10,8 = 559,8 \text{ tới } 602,2 \text{ hạt/phút}$$



Hình 17.1 Phân phối Poisson các giá trị khác nhau của μ . Trục hoành trong các đồ thị là các giá trị của r .

Giả sử có 3 kết quả khác nhau từ buồng đếm lấp lánh dựa trên thời gian khác nhau, như trình bày trong bảng 17.2. Các kết hợp để cho số liệu chung của số đếm trong mỗi phút là cộng 3 số đếm này với nhau và chia cho tổng số thời gian:

$$\text{Số đếm/phút} = 8095/14 = 578,2$$

$$\text{s.e.} = \frac{\sqrt{8095}}{14} = \frac{89,97}{14} = 6,4$$

Số hạt đếm được trong một phút trung bình tương tự như số hạt đếm được trong 5 phút trong thí dụ 17.2. Dù vậy sai số chuẩn nhỏ hơn bởi vì tổng số hạt đếm được trong thời gian dài hơn.

Bảng 17.2 Kết quả từ máy đếm nhấp nháy

Kết quả	Số đếm	Thời gian
1	1740	3
2	2300	4
3	4055	7
Tổng số	8095	14

Phân phối Poisson và tỉ suất

Trong thí dụ 17.2 biến số đo lường là số các phát xạ phóng xạ đếm được trong buồng đếm lấp lánh trong 5 phút nhưng cuối cùng chúng ta tính kết quả là tỉ suất số hạt đếm được trong mỗi phút. Tỉ suất này tính được bằng cách chia số hạt đếm được và sai số chuẩn cho 5. Ở đây chúng ta thảo luận tổng quát cách phân tích tỷ suất và chúng ta dùng kí hiệu λ (kí tự Hi Lạp lambda) cho tỉ suất trung bình số lần xuất hiện biến cố (nghĩa là số lần xuất hiện biến cố trung bình trong một đơn vị thời gian).

Trung bệnh sâu lấu xuất hiều biều cẩu trong mẩut ăn về thài gian $= \frac{\mu}{t} = \lambda$

$$s.e. = \frac{\sqrt{\mu}}{t} = \sqrt{\frac{\lambda}{t}}$$

Trong thí dụ 17.2, ước lượng của μ là 2905, t bằng 5 và λ bằng 581 số hạt đếm được trong mỗi phút. Sai số chuẩn được tính bằng $\sqrt{\mu/t}$ theo cách trên hay $\sqrt{(\lambda/t)}$:

$$s.e. = \sqrt{\frac{581}{5}} = 10,8$$

Cũng giống như trên. Dù vậy công thức thứ hai làm rõ thêm khi thời gian quan sát kéo dài thì sai số chuẩn sẽ giảm, bởi vì bản thân λ cũng như nhau trong thời gian dài và thời gian ngắn.

Phân tích tỉ suất mới mắc

Tỉ suất mới mắc là một loại tỉ suất theo đơn vị thời gian. Nhớ lại từ Chương 15 rằng tỉ suất này bằng số r các trường hợp bệnh mới trong một khoảng thời gian chia cho số người năm nguy cơ. Phân phối Poisson (và xấp xỉ bình thường của nó) có thể được dùng khi có thể giả thiết rằng số các trường hợp xảy ra độc lập với nhau và ngẫu nhiên theo thời gian. Điều này dĩ nhiên ít đúng trong trường hợp bệnh truyền nhiễm hơn là bệnh không truyền nhiễm, nhưng với điều kiện không có bằng chứng mạnh mẽ về sự tập trung của bệnh, việc sử dụng vẫn xứng đáng.

$$\lambda = \frac{r}{\text{người năm nguy cơ}}$$

$$s.e. = \frac{\sqrt{r}}{\text{người năm nguy cơ}}$$

Thí dụ 17.3

Người ta ghi nhận được 47 trường hợp nhiễm trùng hô hấp dưới trong một nghiên cứu 2 năm ở trẻ em dưới 5 tuổi trong một cộng đồng ở Guatemala. Lúc đầu năm trăm trẻ tham gia vào nghiên cứu nhưng bởi vì có sự di chuyển, mới sinh, vượt quá 5 tuổi và không theo dõi được con số quan sát thay đổi theo thời gian. Tổng số trẻ $\times$ năm quan sát được theo dõi là 873. Tỉ suất mới mắc của nhiễm trùng hô hấp dưới, tính trên mỗi 1000 trẻ $\times$ năm được ước tính bằng

$$\lambda = 47/873 \times 1000 = 53,9 \text{ cho mỗi } 1000 \text{ trẻ } \times \text{ năm}$$

Khoảng tin cậy 95% bằng

$$53,9 \pm 1,96 \times 8,6 = 44,3 \text{ tới } 63,5 \text{ nhiễm trùng cho mỗi } 1000 \text{ trẻ } \times \text{ năm}$$

Bảng 17.3 Mới mắc nhiễm trùng hô hấp dưới trong số trẻ dưới 5 tuổi, theo điều kiện gia đình.

Tình trạng gia đình	số nhiễm trùng	trẻ $\times$ năm nguy cơ	tỉ suất mới mắc/1000 trẻ $\times$ năm
Nghèo	33	355	93,0
Tốt	24	518	46,3
Chung	57	873	65,3

Có thể kiểm định bình thường để so sánh hai tỉ suất mới mắc, $\lambda_1 = r_1/\text{người} \times \text{năm}_1$ và $\lambda_2 = r_2/\text{người} \times \text{năm}_2$. Công thức có hiệu chỉnh tính liên tục là:

$$z = \frac{|\lambda_1 - \lambda_2| - \left[\frac{1}{2 \times (\text{ngẫu} \times \text{năm})_1} + \frac{1}{2 \times (\text{ngẫu} \times \text{năm})_2} \right]}{\sqrt{\lambda \left[\frac{1}{(\text{ngẫu} \times \text{năm})_1} + \frac{1}{(\text{ngẫu} \times \text{năm})_2} \right]}}$$

$$\lambda = \frac{r_1 + r_2}{(\text{ngẫu} \times \text{năm})_1 + (\text{ngẫu} \times \text{năm})_2}$$

Trong đó λ là tỉ suất mới mắc toàn bộ trong 2 nhóm và $\lambda_1 - \lambda_2$ là trị số tuyệt đối của hiệu số.

Thí dụ 17.4

Trẻ em được phân tích ở thí dụ 17.3 được chia làm 2 nhóm, sống ở nhà tốt và ở nhà nghèo. Kết quả được trình bày trong bảng 17.3. Tỉ suất mới mắc của nhiễm trùng hô hấp dưới cao hơn đáng kể trong trẻ có điều kiện sống nghèo nàn so với trẻ có điều kiện sống tốt; 93,0 và 46,3 trong 1000 trẻ-năm. Bởi vì tỉ suất này tính theo 100 trẻ-năm, thành phần trẻ năm trong công thức phải được tính theo đơn vị 1000:

$$z = \frac{|93,0 - 46,3| - \left(\frac{1}{1,036} + \frac{1}{0,710} \right)}{\sqrt{75,3 \times \left(\frac{1}{0,518} + \frac{1}{0,355} \right)}} = 2,52$$

Sự khác biệt này có ý nghĩa cao ($P < 0,01$).

TÍNH PHÙ HỢP CỦA PHÂN PHỐI TẦN SUẤT

Giới thiệu

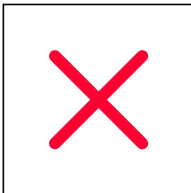
Chúng ta đã gặp một số phân phối lý thuyết. Trong chương này chúng ta sẽ thảo luận cách đánh giá xem phân phối của một nhóm số liệu có tuân theo một mô hình lý thuyết đặc biệt nào đó hay không. Thí dụ, có phải số liệu phù hợp với phân phối bình thường, điều kiện cần cho nhiều phương pháp thống kê. Chúng ta xem phân phối này trước tiên. Phần thứ hai tổng quát hơn. Nó mô tả kiểm định chi bình phương phù hợp (chi square goodness of fit test) để kiểm định ý nghĩa của sự khác nhau giữa phân phối tần suất quan sát được và phân phối được tiên đoán bởi mô hình lý thuyết.

Phù hợp theo phân phối bình thường

Giả thiết về phân phối bình thường là nền tảng cho nhiều phương pháp thống kê. Điều này có thể được kiểm tra bằng cách so sánh hình dạng của phân phối tần suất quan sát được với phân phối bình thường. Ít khi cần đến kiểm định ý nghĩa hình thức (formal) bởi vì chúng ta chỉ quan tâm đến việc khám phá sự sai lệch khỏi tính bình thường một cách đáng kể; hầu hết các phương pháp đều vững vàng đối với sự di lệch vừa phải. Nếu mẫu lớn, việc đánh giá bằng mắt thường là đủ. Thí dụ, hình dạng của tổ chức đồ ở hình 18.1(a) dường như tuân theo phân phối bình thường, trong khi ở hình 18.1(b) rõ ràng bị lệch dương.

Kỹ thuật đồ họa có thể được dùng cho mẫu có kích thước bất kỳ là đồ thị probit (probit plot). Đó là phân tán đồ so sánh điểm phần trăm của phân phối tần suất quan sát với điểm tương ứng (được gọi là probit) của phân phối bình thường chuẩn. Đồ thị probit tuyến tính nếu số liệu phân phối bình thường và cong nếu số liệu không phân phối bình thường.

Thí dụ, bảng 18.1 cho thấy phân phối tần suất của nồng độ hemoglobin của 70 phụ nữ, minh họa ở hình 18.1(a). Không có phụ nữ nào có nồng độ hemoglobin dưới 8 g%. Chỉ có một, 1,4% trong tổng số, có nồng độ dưới 9 g%. Giá trị của phân phối bình thường chuẩn tương ứng với phần trăm này có thể tìm ở bảng A6. Nó bằng -2,20, được gọi là probit của 1,4%. Nhiều người cộng 5 vào giá trị tính được để đảm bảo probit dương nhưng điều này không cần



thiết và sẽ không được tính ở đây.

Ba người phụ nữ có nồng độ hemoglobin từ 9 đến 10 g% vì vậy tổng cộng có 4 người (5,7%) có nồng độ dưới 10 g%. Dùng bảng 6 lần nữa probit tương ứng với 10g% được tính là -1,58. Tương tự có 18 phụ nữ (25,7%) có nồng độ hemoglobin dưới 11 g%. Số được tích lũy theo kiểu này được gọi là số lũy tích (cumulative). Giá trị lũy tích còn lại của phân phối tần suất được trình bày trong bảng 18.1 cùng với probit tương ứng. Người ta không thống nhất cách xử trí đối với 0% và 100% và không có cách nào hoàn toàn thuận lợi. Như ở đây chúng ta bỏ qua chúng và sẽ mất rất ít thông tin.

Bảng 18.1 phân phối tần suất nồng độ hemoglobin của 70 phụ nữ và phân phối tích lũy và probit tương ứng

Hemoglobin	số phụ nữ	số tích lũy	tích lũy(%)	probit
<8	0	0	0,0	
8-	1	1	1,4	-2,20
9-	3	4	5,7	-1,58

10-	14	18	25,7	-0,65
11-	19	37	52,9	-0,07
12-	14	51	72,9	-0,61
13-	13	64	91,4	1,37
14-	5	69	98,6	2,20
15-	1	70	100,0	
16+	0			

Hình 18.1(c) trình bày đồ thị probit và nồng độ hemoglobin. Lưu ý rằng probit -2,20 tương ứng với 9 g% (chứ không phải 8 g%) bởi vì nó tương ứng với phần trăm phân phối nằm dưới 9 g%. Đồ thị này tuyến tính xác nhận rằng số liệu phân phối bình thường. Ngược lại hình 18.1(d) cho thấy đồ thị probit phi tuyến đối với phân phối chiều dày lớp mỡ dưới da vùng cơ tam đầu bị lệch dương.

Đồ thị probit có thể vẽ bằng cách dùng từng quan sát riêng lẻ chứ không dùng phân phối tần suất. Điều này có thể dùng cho mẫu nhỏ. Quan sát được sắp theo thứ tự tăng dần. Phần trăm tích lũy của chúng là:

$$100 \times 1/n, 100 \times 2/n, 100 \times 3/n, \dots, 100 \times n/n$$

Trong đó n là cỡ mẫu, thí dụ nếu có 24 quan sát, phần trăm lũy tích sẽ là 4,2%, 8,4%, 13,2%, ..., 100%. Các probit được tính như trên vào được vẽ theo từng giá trị cá nhân.

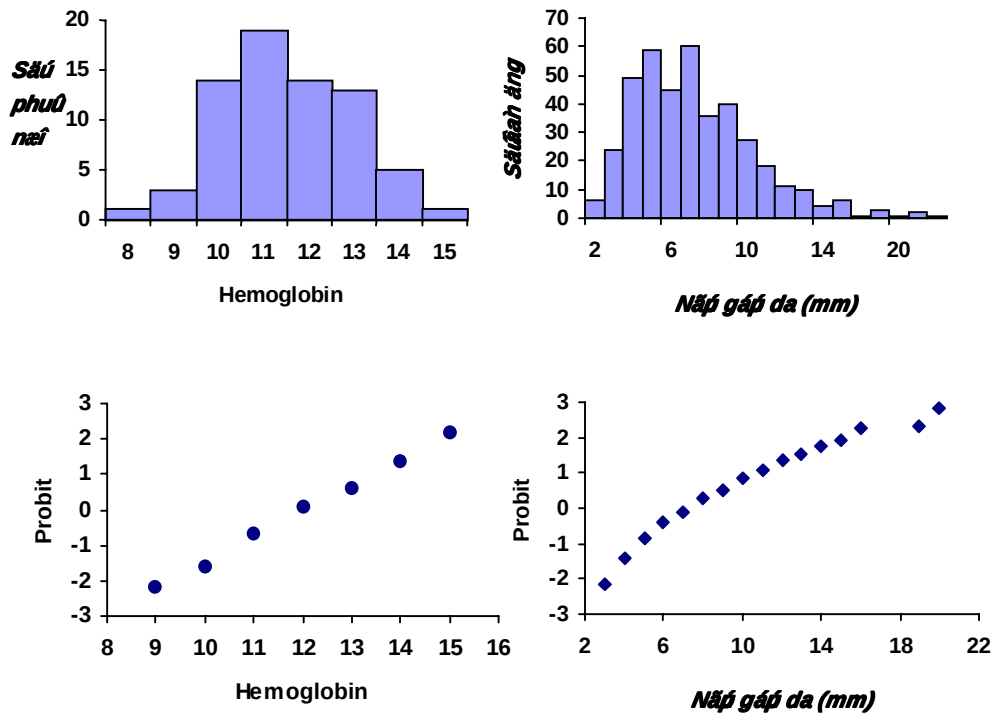
Kiểm định phù hợp chi bình phương

Đôi khi cần kiểm định xem một phân phối tần suất có khác biệt một cách có ý nghĩa với phân phối tần suất lí thuyết giả thiết, như là phân phối Poisson. Có thể tính được bằng cách so sánh tần suất kì vọng và tần suất quan sát dùng kiểm định chi bình phương. Loại kiểm định này cũng giống như trong bảng dự trù, đó là:

$$\chi^2 = \sum \frac{(O - E)^2}{E}$$

Nhưng cách tính độ tự do có khác. Chúng bằng với số nhóm trong phân phối tần suất trừ 1, trừ số thông số tính được từ số liệu. Thí dụ nếu, nếu phân phối mũ phù hợp với thời gian sống thì chỉ cần ước lượng một thông số, đó là 1, nghịch đảo của thời gian sống trung bình. Để phù hợp với phân phối bình thường, cần hai tham số, đó là trung bình m và độ lệch chuẩn s . Trong một số trường hợp không cần ước tính tham số, hoặc là bởi vì mô hình lí thuyết không cần tham số như trong thí dụ 18.1 hoặc là bởi vì các tham số được xác định trong mô hình.

$$d.f. = \frac{\text{Số nhóm trong phân phối tần suất}}{\text{Số các tham số ước lượng}} - 1$$



Hình 18.1 Phân phối tần suất và đồ thị probit để đánh giá tính bình thường của số liệu. (a) và (c) Nồng độ hemoglobin của 70 phụ nữ. Phân phối bình thường. Đồ thị probit tuyến tính. (b) (d) Chiều dày lớp mỡ dưới da vùng cơ tam đầu ở 440 người nam. Lệch dương. Đồ thị probit phi tuyến

Tính toán số kì vọng

Bước đầu tiên trong việc tiến hành kiểm định tính phù hợp chi bình phương là ước lượng tham số cần thiết cho phân phối tần suất lý thuyết từ số liệu. Bước thứ hai là tính số kì vọng trong mỗi nhóm của phân phối tần suất, bằng cách nhân tổng số tần suất cho xác suất một cá nhân giá trị rơi vào trong nhóm

$$\text{Tần suất kỳ vọng} = \text{tần suất tổng cộng} \times \frac{\text{xác suất mẫu cá nhân}}{\text{nằm trong nhóm}}$$

Đối với số liệu rời rạc, xác suất được tính bằng ứng dụng trực tiếp công thức phân phối, như được minh họa trong phân phối Poisson ở thí dụ 18.2. Đối với số liệu liên tục, tính toán từ phân phối tích lũy. Thí dụ, nếu số liệu là thời gian sống còn của muỗi thì xác suất muỗi chết trong ngày thứ ba bằng xác suất nó sống tới đầu ngày thứ ba trừ cho xác suất nó sống tới đầu ngày thứ bốn.

Tính hợp lệ

Kiểm định tính phù hợp chi bình phương không được dùng nếu có một tần suất kì vọng nào nhỏ hơn hai hay có nhiều tần suất kì vọng nhỏ hơn 5. Có thể tránh được bằng cách kết hợp những nhóm kề nhau trong phân phối.

Thí dụ 18.1

Bảng 18.2 trình bày phân phối của kí số cuối cùng (final digit) trong trọng lượng được ghi nhận trong một cuộc điều tra để kiểm tra tính đúng của nó. Chín mươi sáu người lớn được cân và ghi trọng lượng tới 0,1 kg gần nhất. Nếu không có sai lệch trong ghi nhận, thí dụ như sai lệch do khuynh hướng ghi tròn kilogram hay ghi tròn nửa kilogram, ta kì vọng rằng số lần xuất hiện các số 0, 1, 2,..., 9 trong chữ số cuối cùng bằng nhau và bằng 9,6. Có thể kiểm định tính phù hợp của phân phối quan sát được với giả thuyết bằng cách sử dụng kiểm định tính

phù hợp chi bình phương. Có 10 phân phối tần suất và không có tham số nào được ước lượng.

$$\chi^2 = \sum \frac{(O-E)^2}{E} = 9,84 \quad d.f. = 10 - 1 = 9$$

9,84 ở giữa điểm phần trăm 25% và 50% đối với phân phối c2 9 độ tự do. Do đó tần suất quan sát phù hợp với tần suất lí thuyết, gợi ý rằng không có sai lệch ghi nhận.

Bảng 18.2 Kiểm tra tính đúng của cuộc điều tra ghi nhận trọng lượng

Kí số cuối cùng của trọng lượng	tần suất quan sát	tần suất lí thuyết	(O-E) ² /E
0	13	9,6	1,20
1	8	9,6	0,27
2	10	9,6	0,02
3	9	9,6	0,04
4	10	9,6	0,02
5	14	9,6	2,02
6	5	9,6	2,20
7	12	9,6	0,60
8	11	9,6	0,20
9	4	9,6	3,27
Tổng số	96	96,0	9,84

Thí dụ 18.2

Người ta cho rằng sự xâm nhập bội (multiple invasion) của *Plasmodium falciparum* vào hồng cầu xảy ra nhiều hơn so với các kí sinh trùng sốt rét khác. Bảng 18.3 cho thấy số liệu được thu thập để xem những kí sinh trùng này có tấn công vào hồng cầu một cách chọn lọc chửa chỉ đơn giản là sự tình cờ. Năm mươi ngàn hồng cầu được đếm trong một phết máu của bệnh nhân bị sốt rét *falciparum* và ghi nhận số kí sinh trùng trong mỗi hồng cầu

Bảng 18.3 Phân phối tần suất quan sát của số các kí sinh trùng cho mỗi hồng cầu trong máu của bệnh nhân bị *Plasmodium falciparum* so sánh với phân phối Poisson có cùng trung bình. Được pháp của Wang (1970) Transaction of the Royal Society of Tropical Medicine and Hygiene 64, 268-270.

Số kí sinh trùng trong mỗi hồng cầu	số quan sát	phân phối Poisson	(O- E) ² /E
0	40000	39726,7	1,9
1	88621	9137,1	29,2
2	1259	1050,8	41,3
3	99	80,6	4,2
4	21	4,6	68,5
5+	0	0,2	0,2
Tổng số	50000	50000	135,3

Nếu nhiễm trùng một hồng cầu là một biến cố tình cờ, sự phân phối của số các kí sinh trùng cho mỗi tế bào sẽ là phân phối Poisson. Có thể kiểm định giả thuyết này bằng cách dùng kiểm định phù hợp chi bình phương. Bước đầu tiên là tính, μ , số trung bình của kí sinh trùng trong mỗi hồng cầu:

$$\frac{0 \times 40000 + 1 \times 88621 + 2 \times 1259 + 3 \times 99 + 4 \times 21}{50000} = \frac{11520}{50000} = 0,23$$

Bước tiếp theo là tính số kì vọng của hồng cầu có 0, 1, 2, 3, 4, 5+ kí sinh trùng dùng phân phối Poisson với trung bình này. Kết quả được trình bày trong bảng 18.3. Thí dụ, xác suất hồng cầu có hai kí sinh trùng là:

$$\frac{e^{-\mu} \mu^2}{2!} = \frac{e^{-0,23} 0,23^2}{2!} = 0,021015$$

Số kì vọng của hồng cầu trong tổng số 50000 hồng cầu có hai kí sinh trùng là

$$50000 \times 0,021015 = 1050,8$$

Cuối cùng tần suất quan sát và tần suất kì vọng được so sánh với nhau. Có thể thấy rằng có số hồng cầu quan sát có một hay hai kí sinh trùng nhỏ hơn số kì vọng và số hồng cầu quan sát có ba hay bốn kí sinh trùng nhiều hơn kì vọng, gợi ý rằng kí sinh trùng tấn công vào tế bào có chọn lọc chứ không phải ngẫu nhiên. Giá trị chi bình phương có ý nghĩa cao.

$$\chi^2 = 135,3; \text{ d.f.} = 6 - 1 - 1 = 4; P < 0,001$$

PHÉP BIẾN ĐỔI

Giới thiệu

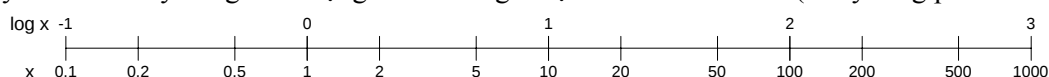
Giả thuyết làm nền tảng cho phương pháp thống kê có thể không phải luôn luôn thỏa mãn bởi một nhóm số liệu nhất định. Thí dụ, phân phối có thể bị lệch dương chứ không phải bình thường, sai số chuẩn của hai nhóm có thể khác nhau không cho phép sử dụng kiểm định t để so sánh hai trung bình, hay quan hệ của hai biến có thể là đường cong chứ không phải là đường thẳng. Ở đây chúng ta sẽ mô tả cách thức vượt qua những vấn đề trên bằng phép biến đổi tới thang đo khác. Sự lựa chọn phổ biến nhất là phép biến đổi logarithm và được mô tả chi tiết. Tóm tắt sử dụng các phép biến đổi sẽ được trình bày ở chương cuối.

Phép biến đổi logarithm

Khi áp dụng phép biến đổi logarithm cho một biến số, mỗi giá trị sẽ được thay thế bằng logarithm của nó

$$u = \log x$$

Trong đó x là giá trị nguyên thủy và u là giá trị biến đổi. Chỉ có thể dùng phép biến đổi này với giá trị dương bởi vì không có logarithm cho giá trị âm và logarithm cho giá trị zero là vô hạn âm. Dầu vậy có một số trường hợp khi cần biến đổi logarithm như trong trường hợp đếm số kí sinh trùng, nhưng số liệu có một số giá trị zero cùng với các giá trị dương. Có thể giải quyết vấn đề này bằng cách cộng 1 vào các giá trị trước khi biến đổi (lưu ý rằng phải trừ đi 1



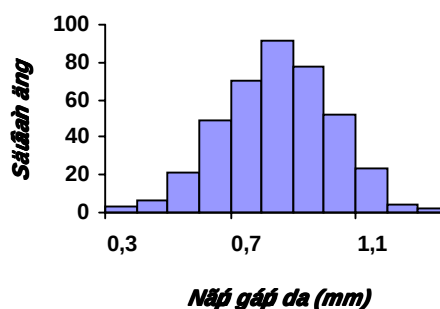
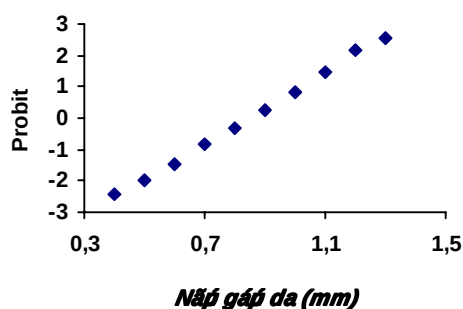
sau khi kết quả cuối cùng được biến đổi về thang đo gốc).

Phép biến đổi logarithm có tác dụng kéo dài phần bên trái của thang đo và nén phần bên phải của thang đo như chỉ ra trong hình 19.1. Thí dụ trong một thang đo logarithm thì khoảng cách giữa 1 và 10 sẽ bằng khoảng cách từ 10 đến 100 và bằng khoảng cách từ 100 đến 1000; chúng cùng có khác biệt gấp 10 lần.

Hình 19.1 Phép biến đổi logarithm

Các giá trị lệch dương

Phép biến đổi logarithm sẽ bình thường hóa các phân phối bị lệch dương, như trong hình 19.2, trong đó là kết quả sự áp dụng phép biến đổi logarithm cho số liệu lớp mỡ dưới da vùng cơ tam đầu được trình bày trong hình 18.1(b). Tổ chức đồ đối xứng và đồ thị probit tuyến tính cho thấy phép biến đổi đã loại bỏ sự lệch và bình thường hóa số liệu. Bề dày lớp mỡ dưới da được gọi là có phân phối log bình thường (lognormal distribution).



Hình 19.2 Phân phối log bình thường của bề dày lớp mỡ dưới da vùng cơ tam đầu của 440 đàn ông. So sánh với hình 18.1

Độ lệch chuẩn không bằng nhau

Cơ chế sử dụng phép biến đổi logarithm sẽ được mô tả bởi việc xem xét số liệu của bảng 19.1(a) cho thấy sự bài tiết β -thromboglobulin (b-TG) ở 12 bệnh nhân tiểu đường cao hơn 12 người bình thường. Những số trung bình này không thể so sánh bằng kiểm định t bởi vì độ lệch chuẩn của hai nhóm khác nhau. Cột bên phải của bảng cho thấy số quan sát sau khi biến đổi logarithm. Thí dụ $\log(4,1)=0,61$. Có thể dùng logarithm với cơ số bất kì; kết quả cũng giống nhau. Người ta thường chọn dùng cơ số 10, như ở đây, hay cơ số e, gọi là logarithm tự nhiên (natural logarithms).

Phép biến đổi logarithm có tác động cân bằng độ lệch chuẩn ở đây (chúng là 0,26 và 0,28 trong thang đo logarithm) và loại bỏ độ lệch ở mỗi nhóm (xem hình 19.3). Kiểm định t có thể được dùng cho thấy log β -TG trung bình cao hơn ở bệnh nhân đái đường một cách có ý nghĩa. Chi tiết tính toán được trình bày ở bảng 19.1(b).

Trung bình nhân và khoảng tin cậy

Khi sử dụng phép biến đổi, tất cả mọi phân tích được tiến hành trên giá trị đã biến đổi, u. Điều cần lưu ý là điều này cũng đúng trong khi tính khoảng tin cậy. Thí dụ, log β -TG trung bình của người bình thường là 1,06 log ng/ngày/100 ml. Khoảng tin cậy của nó là

$$1,06 \pm 2,20 \times 0,26/12 = 0,89 \text{ tới } 1,23 \text{ log(ng/ngày/100ml)}$$

(lưu ý rằng 2,20 là điểm 5% của phân phối t với 11 độ tự do.)

Dù vậy khi báo cáo kết quả cuối cùng đôi khi cần biến đổi ngược thành đơn vị gốc bằng cách lấy antilog như làm trong bảng 19.1(c). Antilog của trung bình của giá trị biến đổi được gọi là trung bình nhân (geometric mean).

$$\text{Trung bình nhân (Geometric mean - GM)} = \text{antilog}(\bar{u}) = 10^{\bar{u}}$$

Bảng 19.1 So sánh β -thromboglobulin (b-TG) nước tiểu ở 12 người bình thường và 12 người bị tiểu đường. Sửa đổi từ kết quả của van Oost et al. (1983). Thrombosis and Haemostasis 49, 18-20, có xin phép.

(a) Số liệu gốc và log của số liệu

β -TG (ng/ngày/100ml creatinine)		Log β -TG (log ng/ngày/100ml creatinine)	
Bình Thường	Tiểu đường	Bình Thường	Tiểu đường
4.1	11.5	0.61	1.06
6.3	12.1	0.80	1.08
7.8	16.1	0.89	1.21
8.5	17.8	0.93	1.25
8.9	24.0	0.95	1.38
10.4	28.8	1.02	1.46
11.5	33.9	1.06	1.53
12.0	40.7	1.08	1.61
13.8	51.3	1.14	1.71
17.6	56.2	1.25	1.75
24.3	61.7	1.39	1.79
37.2	69.2	1.57	1.84
trung bình 13.5	35.3	1.06	1.47
S.D. 9.2	20.3	0.26	0.28
n 12	12	12	12

(b) tính kiểm định t dựa trên log số liệu

$$s = \sqrt{[11 \times 0,262 + 11 \times 0,282/22]} = 0,27$$

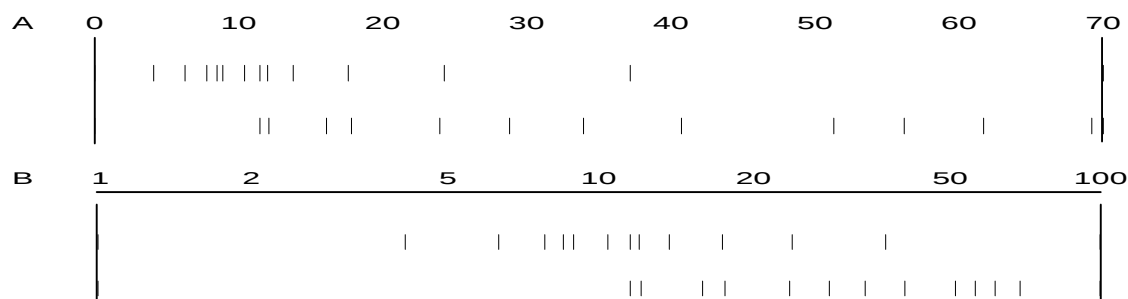
$$t = \frac{1,06 - 1,47}{0,27\sqrt{1/12 + 1/12}} = -3,72 \quad d.f. = 22 \quad P \approx 0,001$$

(c) kết quả báo cáo trên thang đo nguyên thủy

	trung bình nhân		
	β -TG	khoảng tin cậy 95%	độ lệch chuẩn hình học
Bình thường	11,5	7,8 ; 17,0	1,8
Tiểu đường	29,5	19,5 ; 44,5	1,9

Thí dụ, trung bình nhân của β -GT của đối tượng bình thường là:

$$\text{Antilog}(1,06) = 101,06 = 11,5 \text{ ng/ngày/100 ml}$$



Hình 19.3 β -Thromboglobulin data (Bảng 19.1) vẽ từ (a) dùng thang đo tuyến tính và (b) dùng thang đo logarithm. Lưu ý trên thang đo logarithm vẫn ghi theo đơn vị nguyên thủy

Trung bình nhân luôn luôn nhỏ hơn trung bình cộng tương ứng trừ khi tất cả các quan sát có cùng giá trị, trong trường hợp đó hai số đo bằng nhau. Không giống như trung bình cộng, nó không bị ảnh hưởng rõ rệt của các giá trị rất lớn trong phân phối lệch, cho nên nó đại diện cho số trung bình tốt hơn trong tình huống này.

Khoảng tin cậy được tính bằng cách antilog giới hạn tin cậy tính được trên thang đo log. Đối với đối tượng bình thường, khoảng tin cậy 95% của trung bình nhân sẽ là:

$$\text{Antilog}(0,89), \text{antilog}(1,23) = 10^{0,89}, 10^{1,23} \\ = 7,8 ; 17,0 \text{ ng/ngày/100 ml}$$

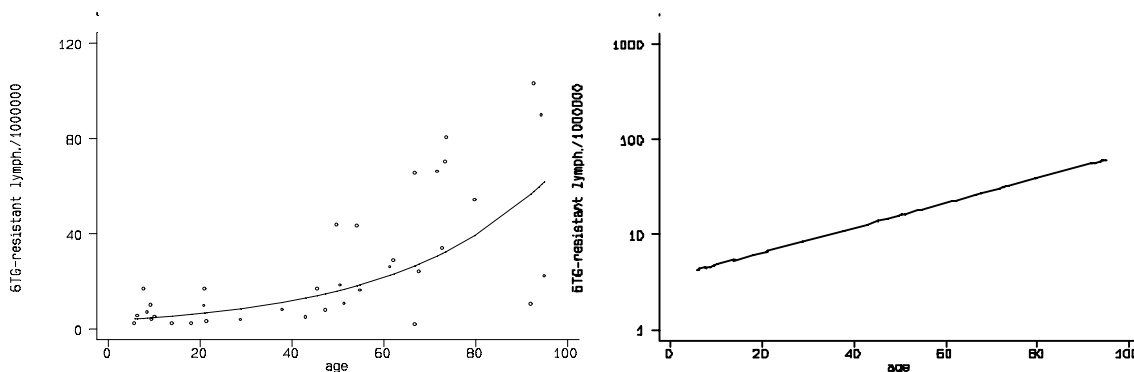
Lưu ý rằng khoảng tin cậy này không đối xứng qua trung bình nhân. Thay vào đó tỉ số giữa giới hạn trên và trung bình nhân, $17,0/11,5 = 1,5$ cũng bằng tỉ số giữa trung bình nhân và giới hạn dưới, $11,5/7,8 = 1,5$. Điều này phản ánh độ lệch chuẩn theo thang đo log tương ứng với sai số nhân chứ không phải sai số cộng trong thang đo gốc.

Vì lí do này, antilog của độ lệch chuẩn khó có thể giải thích và do đó thường ít được dùng. Dù vậy có nhiều tình huống, thí dụ trong bảng lớn có nhiều biến số khác nhau, khi bởi vì tính ngắn gọn, người ta thích dùng antilog của độ lệch chuẩn hơn là khoảng tin cậy. Người ta đã đề nghị từ độ lệch chuẩn hình học (geometric standard deviation) (Kirwood, 1979).

Quan hệ phi tuyến

Hình 19.4(a) trình bày tần suất của lympho bào kháng 6-thioguanine (6TG) gia tăng theo tuổi. Đường cong quan hệ hướng lên và có độ phân tán lớn hơn đối với tuổi lớn hơn. Hình 19.4(b) trình bày bằng cách dùng phép biến đổi log cho tần suất đã làm thẳng mối quan hệ và ổn định độ biến thiên.

Trong thí dụ này, đường cong quan hệ hướng lên và biến số y (tần suất) được biến đổi. Thao tác tương tự cho quan hệ với đường cong đi xuống là lấy logarithm của giá trị x.



Hình 19.4 Quan hệ giữa tần suất của lymphocyte kháng 6TG và tuổi trên 37 đối tượng được trình bày dùng (a) thang đo tuyến tính (b) thang đo logarithm của tần suất. Được phép của Morley *et al.* (1982) *Mechanisms of Ageing and Development* **19**, 21-6

Phân tích hiệu giá

Nhiều thử nghiệm huyết thanh học, như là thử nghiệm kết dính hồng cầu dùng cho kháng thể sốt, dựa trên một loạt những lần pha loãng gấp đôi và độ pha loãng của dung dịch loãng nhất mà có phản ứng được ghi nhận. Kết quả được gọi là hiệu giá và được tính bằng độ pha loãng $1/2, 1/4, 1/8, 1/16, 1/32$ v.v.. Để cho tiện lợi, chúng ta sẽ dùng thuộc ngữ một cách kém chặt chẽ hơn và gọi số nghịch đảo của những số này, đó là 2, 4, 8, 16, 32, v.v. Là hiệu giá. Hiệu giá có khuynh hướng bị lệch dương, và do đó cách phân tích tốt nhất là dùng phép biến đổi logarithm. Điều này được thực hiện đơn giản nhất bằng cách dùng số lần pha loãng thay cho hiệu giá. Do đó hiệu giá 2 được thay bằng số lần pha loãng 1, hiệu giá 4 bằng 2, hiệu giá 8

bằng 3, hiệu giá 16 bằng 4, hiệu giá 32 bằng 5 v.v. Điều này tương đương với lấy logarithm cơ số 2 bởi vì như ta thấy, $8 = 2^3$ và $16 = 2^4$

$u =$ số lần pha loãng $= \log_2$ hiệu giá

Tất cả các phân tích được tiến hành bằng cách dùng số lần pha loãng. Kết quả sau đó được biến đổi ngược trở thành giá trị ban đầu bằng cách tính 2 lũy thừa.

Thí dụ, Bảng 19.2 cho thấy kháng thể sởi của 10 đứa trẻ 1 tháng sau khi tiêm chủng vaccin sởi. Kết quả được tính bằng hiệu giá và số lần pha loãng tương ứng. Trung bình số lần pha loãng là $u = 4,4$. Ta lấy antilog bằng cách tính $2^{4,4} = 21,1$. Kết quả được gọi là hiệu giá trung bình nhân và bằng 21,1.

Hiệu giá trung bình nhân $= 2^{\text{số lần pha loãng trung bình}}$

Bảng 19.2 Nồng độ kháng thể kháng sởi một tháng sau khi tiêm chủng

Trẻ	Hiệu giá kháng thể	số lần pha loãng
1	8	3
2	16	4
3	16	4
4	32	5
5	8	3
6	128	7
7	16	4
8	32	5
9	32	5
10	16	4

Chọn phép biến đổi

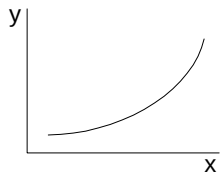
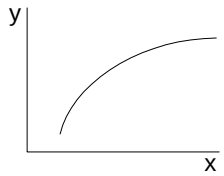
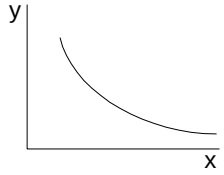
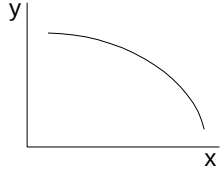
Như đã nói ở trên, phép biến đổi logarithm thường được áp dụng nhiều nhất. Nó thích hợp để loại bỏ tính lệch dương và được dùng trong một số các biến số như thời gian ủ bệnh, thời gian sống còn, số các kí sinh trùng, hiệu giá, liều thuốc, nồng độ chất, và tỉ số. Dù vậy có một số các phép biến đổi cho các số liệu bị lệch được tóm tắt trong bảng 19.3. Thí dụ, phép biến đổi nghịch đảo (reciprocal transformation) mạnh hơn phép biến đổi logarithm và sẽ thích hợp nếu phân phối bị lệch dương nhiều hơn phân phối bình thường, trong khi phép biến đổi lấy căn (root transformation) yếu hơn. Mặt khác tính lệch âm có thể bị loại bỏ bằng cách dùng **phép biến đổi lấy lũy thừa (power transformation) như là biến đổi bậc hai** hoặc bậc ba, độ mạnh tương đương với bậc của lũy thừa.

Cách chọn lựa tương tự cho việc biến đổi độ lệch chuẩn trở nên giống nhau hơn, phụ thuộc vào độ lệch chuẩn tăng như thế nào khi trung bình tăng. (ít khi giảm). Do đó, phép biến đổi logarithm thích hợp nếu độ lệch chuẩn tăng tỉ lệ với độ lệch chuẩn, trong khi phép biến đổi nghịch đảo thích hợp nếu độ dốc cao hơn và phép biến đổi căn thích hợp khi độ dốc ít hơn.

Bảng 19.3 cũng tổng kết một số loại quan hệ phi tuyến có thể xảy ra. Việc chọn lựa phụ thuộc vào hình dạng của đường cong và muốn biến đổi biến x hay biến y .

Cuối cùng ta cũng lưu ý hai phép biến đổi liên quan chặt với tỉ lệ. Chúng là phép biến đổi logistic (hay logit) được giới thiệu ở chương 14 và phép biến đổi probit được mô tả ở chương 18. Tác động chính của hai phép biến đổi này là chuyển một phạm vi giới hạn của thang đo tỉ lệ, từ zero đến 1, thành thang đo vô cực.

Bảng 19.3 Tổng kết sự lựa chọn phép biến đổi. Phép biến đổi làm giảm tính lệch dương gọi là phép biến đổi nhóm A. Phép biến đổi làm giảm tính lệch âm gọi là phép biến đổi nhóm B

Tình huống	Phép biến đổi	
<i>Phân phối lệch dương (nhóm A)</i>		
Log bình thường (lognormal)	logarithm ($u = \log x$)	
Lệch nhiều hơn log bình thường	nghịch đảo ($u = 1/x$)	
Ít lệch hơn phân phối bình thường	căn bậc hai ($u = \sqrt{x}$)	
<i>Phân phối lệch âm (nhóm B)</i>		
Lệch vừa phải	bình phương ($u=x^2$)	
Lệch nhiều	lũy thừa ba ($u = x^3$)	
<i>Biến thiên không đều</i>		
Độ lệch chuẩn tỉ lệ với trung bình	logarithm ($u = \log x$)	
Độ lệch chuẩn tỉ lệ với bình phương của trung bình	nghịch đảo ($u = 1/x$)	
Độ lệch chuẩn tỉ lệ với căn của trung bình	căn bậc hai ($u = \sqrt{x}$)	
<i>Quan hệ phi tuyến</i>	biến y và/hoặc biến x	
	Nhóm A	Nhóm B
	Nhóm B	Nhóm A
	Nhóm A	Nhóm A
	Nhóm B	Nhóm B

PHƯƠNG PHÁP PHI THAM SỐ

Giới thiệu

Phương pháp phi tham số (non-parametric methods) là một nhóm các kỹ thuật thống kê nhằm phân tích số liệu mà không có các giả thuyết về phân phối làm nền tảng, thí dụ như tính bình thường của số liệu. Chúng đặc biệt hữu dụng khi không có tính bình thường trong một nhóm số liệu nhỏ mà không thể sửa chữa bằng phép biến đổi thích hợp.

Nhiều phương pháp phi thống kê được phát minh với mục đích ban đầu là tìm các phương pháp nhanh chóng và dễ dàng nhưng sau đó người ta khám phá ra rằng thành tích của chúng tốt tương đương với những phương pháp tham số cổ điển như kiểm định t. Chúng hầu như có hiệu quả tương đương trong việc phát hiện những sự khác biệt đơn thuần khi giả thuyết tham số được thỏa và tốt hơn đáng kể khi giả thuyết này không thỏa. Như chúng ta sẽ thấy, phương pháp phi tham số rất dễ dùng với điều kiện số liệu không có quá 50 trường hợp. Dù vậy với sự sử dụng rộng rãi máy tính và máy tính cầm tay, điều này trở nên ít quan trọng hơn trước đây. Chúng có hai nhược điểm chính. Thứ nhất, nó chú trọng chủ yếu vào kiểm định ý nghĩa và mặc dù có thể tính được khoảng tin cậy nhưng thao tác phức tạp và tiến hành rất mất công. Thứ nhì, chúng khó mở rộng ra trong các tình huống phức tạp hơn. Vì những lý do này cuốn sách này nhấn mạnh nên dùng phương pháp tham số khi chúng hợp lệ.

Bảng 20.1 Tổng kết những phương pháp phi tham số chính

Phương pháp phi tham số	Sử dụng	Tương đương với phương pháp tham số
Kiểm định dấu (sign test)	Dạng đơn giản của Kiểm định sắp hạng có dấu Wilcoxon	
<i>Kiểm định sắp hạng có dấu Wilcoxon (Wilcoxon sign rank test)</i>	Kiểm định sự khác nhau giữa các cặp quan sát	Kiểm định t cặp đôi
<i>Kiểm định tổng sắp hạng Wilcoxon (Wilcoxon rank sum test)</i>	So sánh 2 nhóm	Kiểm định t hai mẫu
Kiểm định U Mann Whitney (Mann Whitney U test)	Thay thế cho Kiểm định tổng sắp hạng Wilcoxon cho cùng kết quả	Kiểm định t hai mẫu
Kiểm định S Kendall (Kendall's S test)		
Phân tích phương sai một chiều Kruskal Wallis (Kruskal Wallis one-way ANOVA)	So sánh một số nhóm	Phân tích phương sai một chiều
Phân tích phương sai hai chiều Friedman (Friedman two-way ANOVA)	So sánh nhiều nhóm được xác định bởi các giá trị trên 2 biến	Phân tích phương sai hai chiều
<i>Tương quan sắp hạng Spearman (Spearman's rank correlation)</i>	Đo lường sự tương quan giữa hai biến số	Hệ số tương quan
Tương quan sắp hạng Kendall (Kendall's rank correlation)	Thay thế cho tương quan sắp hạng Spearman	Hệ số tương quan
Kiểm định tính phù hợp χ^2 (χ^2 goodness of fit)	So sánh phân phối tần suất quan sát với lý thuyết (xem chương 18)	
Kiểm định Kolmogorov-Smirnov (Kolmogorov-Smirnov test)		
Một mẫu	Thay thế cho kiểm định tính phù hợp χ^2	
Hai mẫu	So sánh hai phân phối tần suất	

Các phương pháp phi tham số chính được liệt kê ở Bảng 20.1 cùng với các phương pháp tham số tương ứng. Có ba phương pháp phổ biến nhất là kiểm định sắp hạng có dấu Wilcoxon, kiểm định tổng sắp hạng Wilcoxon và tương quan sắp hạng Spearman sẽ được mô tả bằng cách dùng cách thí dụ đã được phân tích bởi các kỹ thuật cổ điển. Đối với các kiểm định phi tham số khác độc giả nên đọc Siegel (1956). Nói một cách chặt chẽ, phương pháp χ^2 được mô tả ở chương 13 và 14 cũng là phi tham số nhưng chúng là phương pháp chuẩn cho việc phân tích tỉ lệ và bảng dự trù, nên chúng thường được kể là phương pháp cổ điển.

Kiểm định sắp hạng có dấu Wilcoxon

Đây là kiểm định phi tham số tương đương với kiểm định t cặp đôi. Nó dùng dấu và độ lớn tương đối của số liệu chứ không dùng giá trị thực. Thí dụ xem kết quả ở bảng 20.2 cho thấy số giờ ngủ của 10 người khi dùng thuốc ngủ và khi dùng placebo, và sự khác biệt giữa chúng. Kiểm định sắp hạng có dấu Wilcoxon (Wilcoxon sign rank test) gồm 4 bước căn bản.

Bảng 20.2 Kết quả của thử nghiệm lâm sàng để kiểm định hiệu quả của thuốc ngủ (giống số liệu với Bảng 6.1), cùng với hạng (rank) để dùng cho kiểm định sắp hạng có dấu Wilcoxon.

Bệnh nhân	Số giờ ngủ		hiệu số	hạng (bỏ qua dấu)
	thuốc	placebo		
1	6,1	5,2	0,9	3,5*
2	7,0	7,9	-0,9	3,5*
3	8,2	3,9	4,3	10
4	7,6	4,7	2,9	7
5	6,5	5,3	1,2	5
6	8,4	5,5	3,0	8
7	6,9	4,2	2,7	6
8	6,7	6,1	0,6	2
9	7,4	3,8	3,6	9
10	5,8	6,3	-0,5	1

*Cùng hạng 3 và 4 nên được sắp hạng trung bình

1. Loại bỏ mọi hiệu số bằng zero. Sắp xếp các hiệu số còn lại theo thứ tự tăng dần của độ lớn, bỏ qua dấu và gán cho chúng hạng (rank) 1, 2, 3 v.v.. Nếu hiệu số bằng nhau, như trong trường hợp bệnh nhân 1 và bệnh nhân 2, lấy trung bình hạng của chúng. Các hạng được trình bày trong Bảng 20.2

2. Cộng các hạng có hiệu số dương và các hạng theo hiệu số âm và kí hiệu tổng này là T+ và T-

$$T_+ = 3,5 + 10 + 7 + 5 + 8 + 6 + 2 + 9 = 50,5$$

$$T_- = 3,5 + 1 = 4,5$$

3. Nếu không có sự khác biệt hiệu quả giữa thuốc ngủ và placebo thì tổng T+ và T- phải bằng nhau. Nếu có sự khác biệt thì một tổng sẽ lớn hơn và tổng kia sẽ nhỏ hơn. Kí hiệu tổng nhỏ hơn là T

$$T = \text{số nhỏ hơn của } T+ \text{ và } T-$$

Trong thí dụ này $T = 4,5$

4. Kiểm định sắp hạng có dấu Wilcoxon dựa trên việc đánh giá T, số nhỏ hơn trong hai tổng quan sát, T+ và T-, có nhỏ hơn số nhỏ hơn được kì vọng nếu chỉ có tình cờ. So sánh giá trị thu được với giá trị quyết định (critical value) cho mức ý nghĩa 5%, 2%, 1% được trình bày ở bảng A7. Lưu ý rằng cỡ mẫu thích hợp N là số các hiệu số được sắp hạng chứ không phải tổng số các hiệu số và do đó không kể đến các hiệu số bằng zero.

N = số các hiệu số khác zero.

Ngược lại với các tình huống thường thấy, kết quả có ý nghĩa nếu nó nhỏ hơn giá trị quyết định. Trong thí dụ này cỡ mẫu là mười và các điểm 5%, 2% và 1% tương ứng là 8, 5, 3. Kết quả do đó có mức ý nghĩa 2% bởi vì 4,5 nhỏ hơn 5, chỉ ra rằng thuốc ngủ hiệu quả hơn placebo.

Bảng 20.3 So sánh trọng lượng lúc sinh của con 15 người không hút thuốc lá và trọng lượng lúc sinh của con 14 người hút thuốc lá nặng (số liệu giống như trong Bảng 7.1) với hạng dùng để dùng cho kiểm định tổng sắp hạng Wilcoxon

Không hút thuốc		Hút thuốc lá nặng	
Trọng lượng lúc sinh	Hạng	Trọng lượng lúc sinh	Hạng
3,99	27	3,18	7
3,79	24	2,84	5
3,60*	18	2,90	6
3,73	22	3,27	11
3,21	8	3,85	26
3,60*	18	3,52	14
4,08	28	3,23	9
3,61	20	2,76	4
3,83	25	3,60*	18
3,31	12	3,75	23
4,13	29	3,59	16
3,26	10	3,63	21
3,54	15	2,38	2
3,51	13	2,34	1
2,71	3		
Tổng số = 272		Tổng số = 163	

* Đồng hạng 17, 18, 19 nên được xếp trung bình

Kiểm định tổng sắp hạng Wilcoxon

Đây là kiểm định phi tham số tương đương với kiểm định t. Việc sử dụng kiểm định này được trình bày trong khi xem xét số liệu trong bảng 20.3 trình bày trọng lượng lúc sinh của con 15 người không hút thuốc lá và 14 người hút thuốc lá nặng. Kiểm định tổng sắp hạng Wilcoxon (Wilcoxon rank sum test) bao gồm 3 bước

1. Sắp hạng các quan sát trong cả hai nhóm theo thứ tự độ lớn như trong bảng. Nếu có giá trị nào bằng nhau thì lấy trung bình của chúng

2. Cộng các hạng trong nhóm có cỡ mẫu nhỏ. Trong trường hợp này chính là nhóm người hút thuốc lá và tổng sắp hạng là 163. Nếu hai nhóm có cỡ mẫu bằng nhau thì lấy nhóm nào cũng được.

T = tổng các hạng của nhóm có cỡ mẫu nhỏ

3. So sánh tổng số này với giá trị quyết định trong bảng A8 được sắp đặt hơi khác với các bảng kiểm định ý nghĩa khác. Hãy nhìn vào hàng tương ứng với cỡ mẫu của nhóm, trong trường hợp này nó là hàng 14,15. Phạm vi của mức ý nghĩa 5% là 164 tới 256 và tương ứng với giá trị không có ý nghĩa. Nói cách khác tổng số ở dưới 164 hay trên 256 sẽ là có ý nghĩa ở mức 5%. Tổng 163 trong trường hợp này ở dưới giới hạn dưới 164 có nghĩa là trọng lượng lúc sinh của con những người hút thuốc lá nhỏ hơn con những người không hút thuốc lá một cách có ý nghĩa ($P < 0,05$).

Tương quan sắp hạng Spearman

Đây là số đo phi tham số để đo mức độ liên hệ giữa hai biến số. Phương pháp tham số tương đương là hệ số tương quan đôi khi còn gọi là tương quan tích **moment Pearson (Pearson's product moment correlation)**, được mô tả ở chương 9. Giá trị của mỗi biến số được sắp hạng độc lập và số đo dựa trên hiệu số giữa các cặp số liệu của hai biến. Điều này được minh họa bằng việc sử dụng số liệu trong Bảng 20.4 trong quan hệ giữa thể tích huyết tương và trọng lượng cơ thể của 8 người đàn ông khỏe mạnh. Có 4 bước căn bản:

1. Sắp hạng các giá trị của mỗi biến riêng, như trình bày trong bảng, nếu bất kì giá trị nào giống nhau, lấy trung bình hạng của chúng. Nếu có một số lớn các giá trị giống nhau cần phải dùng công thức điều chỉnh để tính tương quan. Người đọc tham khảo Siegel (1956) để biết chi tiết.

2. Tính hiệu số d, giữa mỗi cặp sắp hạng, bình phương chúng và cộng lại. Trong thí dụ này Sd2 bằng 16

3. Tính tương quan sắp hạng Spearman được kí hiệu bằng r_s

$$r_s = 1 - \frac{6\sum d^2}{n(n^2 - 1)}$$

Trong đó n là số các đối tượng. Tương quan này có giá trị giữa -1 và +1, và cách lí giải tương tự như tương quan thông thường. 1 tương ứng với sự phù hợp hoàn toàn giữa hạng của hai biến số, zero tương ứng với sự không tương quan và -1 tương ứng với sự phù hợp nghịch đảo hoàn toàn giữa các hạng

Trong thí dụ này

$$r_s = 1 - \frac{6 \times 16}{8 \times (8^2 - 1)} = 0,81$$

4. Ý nghĩa của sự liên quan được kiểm định bằng cách so sánh r_s với mức ý nghĩa trong Bảng A9. Giá trị 0,81 dựa trên 8 cặp có ý nghĩa ở mức 5% bởi vì 0,81 lớn hơn giá trị quyết định 0,738 xác nhận quan hệ dương tính giữa thể tích huyết tương và trọng lượng cơ thể.

Khi có hơn 10 cặp, có thể đánh giá ý nghĩa bằng cách khác bằng cách dùng kiểm định t giống như trong hệ số tương quan bình thường.

$$t = r_s \sqrt{\frac{n-2}{1-r_s^2}}, \quad d.f. = n-2$$

Bảng 20.4 Quan hệ giữa thể tích huyết tương và trọng lượng cơ thể ở 8 người đàn ông khỏe mạnh (giống số liệu trong Bảng 9.1) cùng với hạng và hiệu số hạng dùng để tính tương quan sắp hạng Spearman.

Đối tượng	Trọng lượng cơ thể (kg)		Thể tích huyết tương (l)		Hiệu số giữa các hạng	Hiệu số bình phương
	Giá trị(kg)	Hạng	Giá trị (lít)	Hạng	d	d ²
1	58,0	1	2,75	2	-1	1
2	70,0	5	2,86	4	1	1
3	74,0	8	3,37	7	1	1
4	63,5	3	2,76	3	0	0
5	62,0	2	2,62	1	1	1
6	70,5	6	3,49	8	-2	4
7	71,0	7	3,05	5	2	4
8	66,0	4	3,12	6	-2	4
						$\Sigma d^2 = 16$

LẬP KẾ HOẠCH VÀ TIẾN HÀNH NGHIÊN CỨU

Giới thiệu

Từ trước đến giờ chúng ta chỉ nhấn mạnh đến cách phân tích số liệu khi đã thu thập được chúng. Trong 7 chương còn lại chúng ta sẽ thảo luận cách lập kế hoạch (plan) và tiến hành (conduct) một nghiên cứu. Xem xét đầy đủ những vấn đề liên quan có tầm quan trọng hàng đầu, bởi vì dù có phân tích tinh vi đến đâu cũng không thể cứu vớt được một nghiên cứu quy hoạch kém hay được tiến hành sai. Chương này trình bày tổng quan những vấn đề quy hoạch có liên quan, và chương 22 thảo luận về các nguồn sai số có thể phát sinh. Chương 23 tới 25 xem xét các loại quy hoạch nghiên cứu (study design) khác nhau một cách chi tiết và chương 26 xét vấn đề cỡ mẫu, đó là cách tính số các đối tượng cần được nghiên cứu để trả lời câu hỏi quan tâm. Cuối cùng chương 27 giới thiệu về cách sử dụng máy tính.

Không thể trong một phạm vi hạn chế có thể bao trùm tất cả các lãnh vực một cách chi tiết. Mục đích là mô tả đủ để người đọc trở nên quan thuộc với khái niệm liên quan và có thể quy hoạch một nghiên cứu đơn giản và đánh giá các quy hoạch nghiên cứu của các nghiên cứu được đăng trong y văn. Cần tham khảo các sách chuyên môn được liệt kê trong thư mục, thí dụ như về nghiên cứu thử nghiệm lâm sàng và bệnh chứng, để có được những thảo luận chuyên sâu.

Cuối cùng cần lưu ý rằng không có sự đồng ý hoàn toàn về cách phân loại các loại mục tiêu và quy hoạch nghiên cứu khác nhau, cũng chẳng có một ranh giới rõ ràng giữa các loại. Tôi đã chọn và chấp nhận một sơ đồ mà tôi tin là sẽ đạt được một cách tiện lợi, thực tiễn để đánh giá quy hoạch nào thích hợp cho một tình huống nhất định, chứ không phải chú trọng vào các khía cạnh lịch sử của sự phát triển các loại quy hoạch khác nhau.

Mục tiêu của nghiên cứu

Điểm khởi đầu của mọi nghiên cứu là xác định rõ mục tiêu của nó bởi vì chúng sẽ xác định quy hoạch nghiên cứu thích hợp và loại số liệu cần đến. Mục tiêu có thể được phân loại thành một trong ba loại chính dưới đây. Một cuộc nghiên cứu thường có một số mục tiêu, và dĩ nhiên có thể thuộc những loại mục tiêu khác nhau.

- 1. Ước lượng (estimation) một số đặc tính của dân số.** Thí dụ, bao nhiêu phần trăm trẻ em đã được tiêm chủng đầy đủ vaccine vào lúc tròn 3 tuổi hay trung bình số lần bị tiêu chảy hàng năm của trẻ em dưới 5 tuổi ở Bangladesh?
- 2. Nghiên cứu sự liên quan (association) giữa một yếu tố quan tâm và một kết** cuộc nào đó, như là bệnh tật hay là chết. Thí dụ, có phải trẻ em bú mẹ ít bị chết hơn trẻ em không được bú mẹ, hay nhiễm trùng hô hấp phổ biến hơn ở những trẻ mà ba mẹ có hút thuốc lá?
- 3. Lượng giá (evaluation) thuốc hay trị liệu hay một can thiệp nhằm giảm mới** mắc (hay độ nặng) của bệnh. Thí dụ, có phải tiêm BCG giảm nguy cơ bị mụn cơm tái phát, hay việc sử dụng mùng có giảm nguy cơ sốt rét, hay có phải phun thuốc diệt côn trùng vào mùng sẽ giúp bảo vệ thêm?

Phân tích thống kê hệ tịch

Đôi khi có thể trả lời mục tiêu bằng cách dùng thống kê hệ tịch (vital statistics) hay các thống kê y tế thu thập thường quy. Về mặt lịch sử, phân tích thống kê hệ tịch được ghi nhận đầu tiên được tiến hành bởi John Graunt ở Luân đôn, người đã xuất bản cuốn sách *Natural and Political Observations upon the Bills of Mortality* vào năm 1662. ngoài các điều khác, Graunt còn khẳng định rằng nam nhiều hơn nữ trong cả sinh và tử, tỉ suất tử vong trẻ em cao và sự biến thiên theo mùa của tử vong. Năm 1663, Edmund Halley dùng các giấy khai tử để xây dựng bảng sống đầu tiên (xem Chương 16). Graunt đi trước thời đại của ông và truyền thống phân tích không được tiến hành cho đến giữa thế kỉ 19 khi William Farr trở thành Tổng

Đăng kiểm (Registrar General) đầu tiên ở Anh và xứ Wales. Farr đã áp dụng số liệu thống kê hộ tịch vào nhiều vấn đề y tế công cộng bao gồm việc so sánh các tỉ suất tử vong giữa các nhóm nghề nghiệp khác nhau.

Xem xét các số liệu thống kê hộ tịch được thu thập thường quy về sinh, tử vong và bệnh tật cũng gợi ý cho sự liên hệ giữa bệnh tật và nguyên nhân. Thí dụ, sự gia tăng tử vong vì ung thư phổi và sự liên hệ với sự gia tăng tần suất hút thuốc lá đầu tiên được lưu ý trong số liệu thống kê hộ tịch. Giả thuyết này sau đó được kiểm định với các nghiên cứu bệnh chứng và đoàn hệ được quy hoạch đặc biệt (xem ở dưới).

Nghiên cứu quan sát

Nói chung, cần tiến hành một nghiên cứu đặc biệt để thu thập số liệu thích hợp để trả lời cho một mục tiêu chuyên biệt. Mục tiêu ước lượng và quan hệ dẫn tới những nghiên cứu có bản chất quan sát (observational); diễn tiến tự nhiên của bệnh được quan sát mà nghiên cứu không có ý muốn thay đổi nó. Mục tiêu đánh giá có thể được trả lời bởi nghiên cứu quan sát hay thực nghiệm (experimental), phụ thuộc vào loại biện pháp được đánh giá và biện pháp đó đã được sử dụng hay chưa. Một thí dụ của nghiên cứu quan sát là đánh giá hiệu quả của tiêm BCG chống lại bệnh lao ở một vùng có chương trình tiêm chủng tích cực ở đó BCG được tiêm chủng thường quy, bởi vì việc lựa chọn có tiêm hay không tiêm không phụ thuộc vào người nghiên cứu. Ngược lại, thử một loại vaccine chống sốt rét mới có lẽ là thực nghiệm.

Quy hoạch nghiên cứu quan sát có thể chia thành ba loại chính, cắt ngang, cắt dọc (kể cả đoàn hệ) và nghiên cứu bệnh chứng. Sơ đồ lấy mẫu thích hợp cho nghiên cứu cắt ngang và cắt dọc được mô tả chi tiết ở Chương 23. Nói chung mục tiêu là sử dụng phương pháp mà mọi người có cơ hội được chọn bằng nhau, đó là sơ đồ **xác suất cân bằng (equal probability scheme)**. **Người có thể được chọn riêng cá nhân** hay chọn thành cụm, mặc dù cần một cỡ mẫu lớn hơn cho việc chọn cụm. Dù vậy các khó khăn tài nguyên và hậu cần thường làm cho không thể xét một số lượng các cụm đủ để có một bức tranh đại diện và kết quả có khi chỉ dựa trên điều tra một cộng đồng. Đánh giá tính đại diện là vấn đề chủ quan chứ không phải thống kê. Điều này ít quan trọng nếu mục tiêu chính là nghiên cứu sự liên quan chứ không phải là ước lượng. Phương pháp đo lường sự liên quan được trình bày ở Chương 24, ở đó cũng xét các khía cạnh quy hoạch đặc hiệu của nghiên cứu bệnh chứng.

Nghiên cứu cắt ngang

Nghiên cứu cắt ngang (cross-sectional study) được tiến hành tại một thời điểm trong thời gian hay trong một khoảng thời gian ngắn. Nghiên cứu cắt ngang rất nhanh, rẻ, và dễ tiến hành và có thể phân tích ngay. Bởi vì chúng có thể ước lượng những đặc tính của cộng đồng tại một thời điểm chúng thích hợp cho việc đo lường bệnh toàn bộ chứ không thích hợp cho bệnh mới mắc và khó giải thích mối liên quan tìm thấy. Thí dụ, một cuộc điều tra về onchocerca cho thấy có hai khả năng của sự liên hệ này. Thứ nhất là những người có tình trạng dinh dưỡng kém có sức đề kháng kém và do đó dễ bị mù hơn do onchocerca. Thứ nhì là tình trạng dinh dưỡng kém là hậu quả của không phải là nguyên nhân của mù lòa, bởi vì người mù không thể tự kiếm sống. Cần nghiên cứu cắt dọc để quyết định lời giải thích nào là hợp lý hơn.

Nghiên cứu cắt dọc

Trong nghiên cứu cắt dọc (longitudinal study) các cá nhân được theo dõi theo thời gian, khiến cho có thể đo lường bệnh mới mắc và các thay đổi theo thời gian và dễ dàng nghiên cứu diễn tiến tự nhiên của bệnh. Trong một vài tình huống có thể theo dõi số liệu sinh, tử vong, đợt bệnh bằng việc giám sát liên tục (continuous monitoring), thí dụ bằng việc giám sát các hồ sơ đăng kí trong dân số có hệ thống khai tử. Đôi khi việc thu thập số liệu là hồi cứu (retrospective) được tiến hành trong các hồ sơ cũ. Phổ biến hơn nó có thể là tiền cứu (prospective) và vì lí do này, nghiên cứu cắt dọc thường được gọi là nghiên cứu tiền cứu.

Trong đa số các trường hợp, cách đơn giản nhất để tiến hành nghiên cứu cắt dọc là tiến hành nhiều lần điều tra cắt ngang (repeated cross-sectional surveys) ở những khoảng thời gian cố định và điều tra hay đo lường những thay đổi đã xảy ra giữa các cuộc điều tra, như sinh, tử vong, di trú, thay đổi trọng lượng hay mức kháng thể, hay xuất hiện các đợt bệnh mới. Khoảng được chọn phụ thuộc vào các yếu tố được nghiên cứu. Thí dụ, để đo lường mức mắc tiêu chảy, được đặc trưng bởi sự lặp đi lặp lại những cơn bệnh ngắn, số liệu cần thu thập hàng tuần để đảm bảo nhớ lại chính xác. Để giám sát sự phát triển trẻ em chỉ cần đo lường mỗi một hay ba tháng.

Dân số nghiên cứu có thể là động hay cố định. Trong một dân số động (dynamic) các cá nhân rời khỏi nghiên cứu khi nó không còn tuân theo định nghĩa dân số trong khi các cá nhân mới thỏa mãn điều kiện có thể tham gia. Một thí dụ là nghiên cứu mức mắc của bệnh tiêu chảy trong trẻ dưới 5 tuổi, trong đó việc giám sát sẽ dừng lại khi nó đạt được 5 tuổi trong khi các trẻ mới sinh được chiêu mộ vào dân số khi mới sinh. Trong tình huống cố định (fixed), dân số được xác định lúc đầu và ngoài lý do chết, di chuyển hay mất dấu theo dõi, cấu trúc dân số vẫn giữ nguyên trong suốt nghiên cứu. Một dân số cố định được gọi là đoàn hệ (cohort), như là các **đoàn hệ sinh vào năm 1984 (birth cohort of 1984), đó là tất cả những người** cùng sinh năm 1984. Một thí dụ khác là một đoàn hệ công nghiệp những người làm trong công nghiệp năng lượng nguyên tử giữa 1970 và 1974.

Bù lại nhiều ưu điểm, nghiên cứu cắt dọc có một số khuyết điểm. Thứ nhất việc phân tích phức tạp hơn trong nghiên cứu cắt ngang và có thể cần những phương tiện xử lý số liệu tinh vi. Việc liên kết kết quả giữa các cuộc điều tra không dễ dàng và phải chấp nhận sự di chuyển, tử vong hay nhập cuộc mới. Cần điều chỉnh cho việc già đi của nhóm theo thời gian. Thứ nhì, cần thời gian nghiên cứu dài hơn và nhiều cuộc điều tra, tỉ suất bỏ cuộc và không đáp ứng cao hơn và vấn đề số liệu không hoàn chỉnh lớn hơn. Cuối cùng nghiên cứu cắt dọc có thể tốn kém hơn và đặc ra nhiều vấn đề hậu cần trong việc thực hiện.

Nghiên cứu bệnh chứng

Nghiên cứu bệnh chứng (case control study) được dùng để tìm hiểu sự liên hệ giữa một yếu tố nhất định và một bệnh nhất định. Việc quy hoạch rất khác với những loại nghiên cứu khác bởi vì được việc lấy mẫu được tiến hành tùy theo tình trạng bệnh tật chứ không phải theo tình trạng tiếp xúc. Một nhóm các đối tượng được xác định là có bệnh, bệnh (case) được so sánh với một nhóm không có bệnh, **chứng (control)**. **Thí dụ để đánh giá việc bú sữa mẹ có bảo vệ đứa trẻ không bị chết**, nhóm bệnh gồm những đứa trẻ bị chết trong năm đầu tiên và nhóm chứng là những đứa trẻ sống trong cùng khu vực có cùng tuổi và giới tính với bệnh. Nếu giả thuyết đúng thì số bệnh sẽ ít bú sữa mẹ hơn nhóm chứng.

Nghiên cứu bệnh chứng đặc biệt thích hợp với các bệnh hiếm, bởi vì nó cần một cỡ mẫu nhỏ hơn đáng kể so với nghiên cứu cắt ngang hoặc cắt dọc tương ứng. Chúng tương đối rẻ tiền, nhanh chóng, dễ tiến hành và cho phép tính số mới mắc tương đối của bệnh (xem Chương 24) trong nhóm tiếp xúc và nhóm không tiếp xúc với yếu tố quan tâm. Tính phức tạp của phương pháp nằm ở chỗ cần phải xem xét cẩn thận trong việc quy hoạch để giảm thiểu sai lệch (bias) và các phân tích tương đối tinh vi cần thiết. Đặc biệt vấn đề chọn nhóm chứng cũng là vấn đề gây tranh luận. Nghiên cứu bệnh chứng thường thành công nhất trong việc phát hiện những tác động lớn, bởi vì nó dễ dàng thuyết phục rằng tác động đó là thực chứ không phải là hình ảnh giả tạo do quy hoạch. Cuối cùng, nghiên cứu bệnh chứng (giống như nghiên cứu cắt ngang) không thích hợp nếu như bệnh hướng đến yếu tố nguy cơ và cũng có thể là hậu quả của yếu tố nguy cơ như trong trường hợp suy dinh dưỡng là yếu tố nguy cơ của tiêu chảy.

Nghiên cứu thực nghiệm

Nghiên cứu y khoa có bản chất thực nghiệm có thể phân làm 3 nhóm chính tùy theo bản chất của đo lường được lượng giá

1. Thử nghiệm lâm sàng (clinical trials)) thuốc hay các biện pháp điều trị

2. Thử nghiệm vaccine (vaccine trials)

3. Thử nghiệm can thiệp (intervention trials)

(a) các chế độ dự phòng (không phải vaccine), thí dụ cho thuốc chống sốt rét trong 5 năm đầu tiên của cuộc sống cho trẻ ở vùng bị lưu hành sốt rét. Những nghiên cứu này đôi khi còn được gọi là thử nghiệm dự phòng (prophylactic trials)

(b) Các biện pháp phòng ngừa, thí dụ như dùng mùng bảo vệ chống sốt rét.

Việc quy hoạch được thảo luận chi tiết ở Chương 25. Thành phần căn bản là việc định vị cá nhân vào các nhóm thực nghiệm khác nhau dưới sự kiểm soát của người nghiên cứu, và sự định vị này phải đảm bảo các nhóm càng giống nhau càng tốt về các mặt ngoại trừ chế độ hay biện pháp điều trị được lượng giá. Nên sử dụng thử **nghiệm có kiểm soát mù đôi (double-blind controlled trial)** (xem Chương 25) để tránh sai lệch (bias).

Quy hoạch bản vấn lục

Trong hầu hết các cuộc nghiên cứu cần phải chuẩn bị một biểu mẫu ghi nhận được quy hoạch đặc biệt hay một bản vấn lục để thu thập số liệu. Nó phải càng ngắn gọn càng tốt và cần phải tránh sự cảm dỗ muốn hỏi tất cả các câu hỏi. Bảng vấn lục quá dài sẽ làm mệt mỏi người phỏng vấn và người được phỏng vấn và có thể dẫn tới câu trả lời sai lầm. Câu hỏi phải rõ ràng và không mơ hồ và phải viết rõ ràng bởi vì chúng sẽ được đọc. Phải tránh các từ kỹ thuật hay các từ dài cũng như các câu hỏi âm tính, các câu hỏi gợi ý và các câu hỏi giả thuyết.

Cần phải suy nghĩ cẩn thận về thứ tự thu thập thông tin và bảng vấn lục sẽ được tự điền hay được người phỏng vấn hoàn thành. Phải có một tiến trình logic xuyên suốt biểu mẫu để cho nó dễ dàng theo dõi. Các câu hỏi được sắp xếp thành các phần. Cần phải chỉ rõ việc hoàn tất một phần nào đó có phụ thuộc vào câu hỏi trước đó hay không, và điểm tiếp theo của phần phải được nhận dạng dễ dàng. Tốt nhất nên giảm thiểu số bước nhảy có thể xảy ra bởi vì có thể có nhầm lẫn và các phần có thể bị bỏ sót. Bản vấn lục phải được ký hiệu với tên của nghiên cứu và tên người trả lời và số danh bộ nghiên cứu. Nó phải được lập lại trên mỗi trang. Phần tiếp theo thường gồm các thông tin nhận dạng như tuổi, giới tính, và địa chỉ. Nói chung cần sắp xếp các phần tiếp theo theo thứ tự quan trọng của thông tin đối với nghiên cứu, để cho phần thông tin quan trọng nhất được thu thập khi người phỏng vấn và người được phỏng vấn đều thoải mái nhất và ít chán nhất. Dù vậy, bất cứ câu hỏi tế nhị nào phải để sau cùng.

Câu hỏi đóng và câu hỏi mở

Câu hỏi có 2 dạng, mở và đóng. Câu hỏi mở (open question) được dùng để tìm kiếm thông tin và người phỏng vấn ghi nhận câu trả lời trong dạng viết tự do. Không có định kiến về và câu trả lời nào sẽ xảy ra. Mặt khác trong câu hỏi **đóng (close question)**, câu trả lời giới hạn ở trong danh sách các câu trả lời có thể. Danh sách này nên có một mục cho 'các ý kiến khác' với chỗ trống để viết chi tiết và mục 'không biết'. Người phỏng vấn có hướng dẫn người trả lời từ mục này tới mục khác trong danh sách hay hỏi dưới dạng mở và đánh dấu mục tương ứng với câu trả lời. Câu hỏi đóng và câu hỏi mở có ưu điểm và khuyết điểm riêng và việc chọn lựa phụ thuộc rất nhiều vào hoàn cảnh cụ thể. Trả lời cho câu hỏi đóng dễ xử lý nhưng câu hỏi mở cho nhiều chi tiết thông tin sâu sắc hơn. Một khả năng là dùng câu hỏi mở trong giai đoạn nghiên cứu dẫn đường và dùng kết quả để rút ra danh sách các câu trả lời cho câu hỏi đóng trong nghiên cứu chính.

Mã hóa

Các số liệu số phải được ghi nhận càng chi tiết càng tốt và nên ghi nhận giá trị riêng biệt chứ không mã hóa trước trong bản vấn lục thành nhiều nhóm. Xét thí dụ về tuổi. Các lựa chọn tốt là ghi ngày sinh và sau đó tính tuổi từ ngày sinh và ngày phỏng vấn. Lựa chọn tốt tiếp theo là ghi tuổi của người trả lời, thí dụ như năm cho người lớn, tháng cho trẻ nhỏ, tuần cho trẻ nhũ nhi và ngày cho trẻ sơ sinh. Các kém thuận lợi nhất là ghi nhận người trả lời thuộc nhóm tuổi nào như từ 0-4, 5-9, 10-14, 15-24, 25-44, 45-64 và 65+ tuổi. Đơn vị đo lường phải được ghi

rõ, thí dụ trọng lượng được tính theo kilogram hay pound và phải chỉ số số các số lẻ. Với câu hỏi đóng, mã hóa số tương ứng cần phải in dọc theo danh sách các câu trả lời. Với câu hỏi mở, cần phải mã hóa câu trả lời sau khi hoàn thành bản vấn lục và phải chừa trống cho việc mã hóa.

Nếu số liệu được nhập vào máy tính cần phải nhớ điều này khi quy hoạch bản vấn lục. Thí dụ nó có thể làm nhanh quá trình nhập liệu nếu tất cả các thông tin cần nhập được mã hóa thành ô dọc theo bên lề phải của biểu mẫu. Trong đa số trường hợp cần gán mã số cho câu trả lời không phải là số như mã 1 cho đàn ông và 2 cho đàn bà. Cần một ô cho mỗi số (hay mỗi ký tự), tổng số các ô cần thiết cho biến được xác định bằng số tối đa các chữ số cần ghi nhận cho biến đó. Nếu có thể tránh sử dụng mã zero bởi vì nhiều chương trình máy tính và nhiều gói phần mềm thống kê không thể phân biệt zero với không trả lời, không có số liệu.

Câu hỏi nhiều trả lời

Câu hỏi nhiều trả lời (multiple response questions) cần phải xem xét đặc biệt. Nó có thể xử trí bằng 2 cách. Thí dụ, trong nông thôn Tây Phi, một gia đình có thể dùng một hay nhiều trong tám loại nguồn nước (nước mưa, nước giếng sâu, nước giếng, suối phun, sông, hồ, cao, suối) để uống. Phương pháp thứ nhất là gán một ô cho mỗi câu trả lời có thể, trong trường hợp này là 8 ô. Mỗi ô sẽ chứa mã '1' cho việc sử dụng và '2' cho việc không sử dụng nguồn đó. Nếu một gia đình dùng nước mưa và cũng lấy nước từ sông, ô nước mưa (số một) và ô nước sông (số năm) sẽ chứa mã '1' trong khi 6 ô còn lại sẽ chứa mã '2'. Phương pháp thứ nhì là gán mã 1 tới 8 cho 8 loại nguồn nước khác nhau và quyết định số câu trả lời tối đa mà gia đình có thể trả lời. Giả sử rằng chúng ta quyết định giới hạn là 3 trả lời. Chúng ta sẽ để 3 ô mã hóa cho câu hỏi và nhập mã số của nguồn nước. Gia đình dùng nước mưa và nước sông sẽ được mã hóa 1 (nước mưa) trong ô 1 và 5 (nước sông) trong ô 2. Ô 3 sẽ để trống. Mã có thể được nhập theo thứ tự của số, như chúng ta đã làm, hay theo một thứ tự logic, như theo số lượng nước sử dụng của mỗi nguồn.

Kiểm tra số liệu

Mọi bản vấn lục phải được kiểm tra kỹ lưỡng sau khi hoàn tất và khi số liệu được nhập liệu. Không thể bỏ qua tầm quan trọng của vấn đề này. Kiểm tra phải được tiến hành ngay sau khi thu thập số liệu để có cơ hội tối đa giải quyết các thắc mắc. Kiểm tra có hai loại chính, kiểm tra phạm vi và kiểm tra tính phù hợp. Kiểm tra **phạm vi (range checks) loại bỏ, chẳng hạn như xuất hiện mã 3 cho giới tính, mà lẽ ra chỉ là mã 1 (nam) hay mã 2 (nữ)**. Kiểm tra tính phù hợp (consistency checks) phát hiện những tổ hợp không thể có của số liệu như người đàn ông mang thai, hay 3 tuổi nặng 70 kg. Phân tán đồ cho thấy sự liên hệ giữa hai biến đặc biệt hữu ích trong trường hợp này, bởi vì nó cho phép phát hiện dễ dàng các tổ hợp kỳ lạ.

Cần lưu ý 3 điểm cơ bản để giảm thiểu sai số xảy ra trong xử lý số liệu. Thứ nhất là tránh các việc sao chép không cần thiết số liệu từ biểu mẫu này tới biểu mẫu khác. Thứ nhì là dùng thao tác kiểm chứng (verification) trong khi nhập liệu. Số liệu được nhập hai lần, tốt nhất là bởi hai người khác nhau, bởi vì nó cho phép đánh giá độc lập các con số viết thẩu. Hai tập hợp số liệu được so sánh với nhau và giải quyết các điểm không thống nhất. Thứ ba là kiểm tra các tính toán cẩn thận, hoặc là bằng sự lặp lại hay là chẳng hạn như kiểm tra lại tổng từng phần cộng lại có bằng tổng số chung hay không. Khi sử dụng máy tính, cần phải thao tác ban đầu trên một nhóm số liệu nhỏ và kiểm tra lại bằng tính tay.

NGUỒN GỐC SAI SỐ

Giới thiệu

Điều quan trọng là nhận thức rằng các loại sai số khác nhau có thể xâm nhập vào việc quy hoạch, tiến hành và phân tích nghiên cứu để có thể giảm thiểu sự xuất hiện sai số. Chúng ta bắt đầu bằng sự phân biệt giữa hai loại chính, sai số hệ **thống (systematic)** và **ngẫu nhiên (random)**. **Thí dụ, xét kết quả một thực nghiệm** để đánh giá sự phù hợp giữa các điều dưỡng trong việc đo huyết áp. Bốn điều dưỡng cùng đo 10 người tình nguyện. Việc xem xét kết quả cho thấy mặc dù các điều dưỡng A, B, C không có cùng một kết quả như nhau nhưng không có một mô thức không phù hợp và giá trị trung bình giống nhau. Do đó mọi sai số xuất hiện đều là ngẫu nhiên. Ngược lại điều dưỡng D đọc huyết áp luôn luôn thấp hơn các cộng sự của cô và trung bình của cô thấp hơn 10 mmHg. Cô mắc phải sai số hệ **thống** và **ước lượng trung bình của cô ta không đúng, thấp hơn trị số thực**. Kết quả của cô ta bị sai lệch (bias).

Mỗi giai đoạn nghiên cứu đều có thể bị sai số hệ thống. Chúng nguy hiểm bởi vì sai lệch có thể dẫn đến những kết luận không hợp lệ. Mặc khác sai số ngẫu nhiên giảm độ chính xác nhưng không ảnh hưởng tới tính hợp lệ. Chúng xảy ra trong khi thu thập số liệu do bản vấn lục hay các thiết bị bị hỏng, quan sát sai, sai lầm người trả lời và trong quá trình xử lý số liệu do việc mã hóa, sao chép, nhập liệu, viết chương trình, sai số tính toán. Sai số cũng phát sinh trong khi chọn phương pháp phân tích không đúng và lý giải sai kết quả.

Sai lệch có thể được phân loại thành 3 loại chính: sai lệch chọn lựa, sai lệch thông tin và sai lệch gây nhiễu. Tất cả chúng đều dễ xảy ra trong nghiên cứu quan sát hơn trong nghiên cứu thực nghiệm, ở đó chúng có thể tránh được bằng cách dùng quy hoạch mù đôi có kiểm soát ngẫu nhiên hóa (randomized controlled double-blind design) (xem Chương 25). Chúng đặc biệt là vấn đề trong quy hoạch bệnh chứng. Chúng ta sẽ xét ngắn gọn các loại sai lệch. Để có thêm chi tiết tham khảo Kleinbaum et al. (1982). Sau đó chúng ta sẽ xét hai tham số, độ nhạy cảm và độ đặc hiệu, được ung để đánh giá khả năng một thao tác phân loại đúng tình trạng bệnh tật (hay tiếp xúc) của một cá nhân. Cuối cùng, chúng ta sẽ xét một trường hợp đặc biệt của sai lệch, phát sinh từ hiện tượng gọi là hồi quy về trung bình.

Sai số chọn lựa

Sai số chọn lựa (selection bias) có thể do một (hay cả hai) loại khiếm khuyết khác nhau trong quy hoạch về đối tượng được chọn nghiên cứu. Thứ nhất là nếu đối tượng được chọn lọc khác biệt có hệ thống với các đối tượng không được chọn lọc. Nguyên nhân chính là tỉ suất không đáp ứng cao, thất lạc theo dõi hay cơ cấu lấy mẫu sai. Thí dụ, trong nhiều quốc gia việc nghiên cứu các trường hợp tiêu chảy nặng dựa vào bệnh viện sẽ loại bỏ các trường hợp cấp tính chết trước khi đến bệnh viện và sẽ loại bỏ chọn lọc những người sống xa bệnh viện và những gia đình có tình trạng kinh tế xã hội hay giáo dục kém, bởi vì họ là những người ít sử dụng các phương tiện y tế. Loại thứ hai là sai lệch chọn lọc trong nghiên cứu so sánh do việc chọn nhóm so sánh không đúng, khác biệt về mức độ theo dõi y khoa, về việc nhập viện, dẫn tới khả năng được chẩn đoán trong nhóm tiếp xúc cao hơn so với nhóm không tiếp xúc (ngộ biến Berkson - Berkson's fallacy), khác biệt sống còn hay khác biệt theo dõi.

Sai lệch gây nhiễu

Sai lệch gây nhiễu (confounding bias) xảy ra khi có sự khác biệt quan trọng giữa các nhóm được so sánh mà sự khác biệt này có liên quan đến biến số khác tâm. Một thí dụ đã được thảo luận ở Chương 14 trong Bảng 14.4, trong đó chúng ta đã so sánh mắc toàn bộ của leptospira trong người lớn ở nông thôn và thành thị. Trong trường hợp này giới tính là biến số gây nhiễu, bởi vì không chỉ tỉ suất mắc toàn bộ cao hơn ở nam so với nữ mà cấu trúc dân số của thành thị và nông thôn khác nhau. chúng ta có thể khắc phục bằng cách xây dựng những bảng 2×2 riêng cho nam và cho nữ và sử dụng kiểm định χ^2 Mantel-Haenszel.

Vấn đề gây nhiễu là một vấn đề có thể xảy ra cho mọi nghiên cứu. Có thể kiểm soát bằng quy hoạch sử dụng kỹ thuật bắt cặp mô tả ở Chương 24 và 25, và trong nghiên cứu thực nghiệm bằng cách ngẫu nhiên hóa các cá nhân vào các nhóm thực nghiệm khác nhau. Không giống như các loại sai lệch khác, sai lệch gây nhiễu có thể được điều chỉnh trong giai đoạn phân tích.

Sai lệch thông tin

Sai lệch thông tin (information bias) do sự đo lường hay trả lời sai một cách có hệ thống hay từ việc phân loại sai các tình trạng tiếp xúc và bệnh tật. Nó có nhiều nguyên nhân khác nhau. Các nguyên nhân bao gồm các sai sót trong bản vấn lục, các sai số quan sát, các sai số do người trả lời, sai số do dụng cụ. Sai sót trong bản vấn lục có thể do câu hỏi không đúng về mặt văn hóa, cách dùng từ mơ hồ hay có quá nhiều câu hỏi. Sai số quan sát có thể do hiểu lầm các thao tác, lí giải sai các câu trả lời, hay chỉ do sai lầm. Nghiêm trọng hơn, các sai số quan sát có thể làm sai lệch phát hiện của nghiên cứu do sự khác biệt trong đánh giá, trong thăm dò hay trong việc lí giải các câu trả lời còn nghi ngờ, nếu người quan sát biết được tình trạng tiếp xúc hay bệnh tật của người trả lời. Điều này dễ xảy ra khi người phỏng vấn biết giá thuyết của nghiên cứu.

Sai số do người trả lời có thể bắt nguồn từ việc hiểu lầm, nhớ lại sai lầm, trả lời bằng các câu trả lời được cảm nhận là 'đúng' hay do thiếu quan tâm. Trong một số trường hợp người trả lời có thể trả lời sai có ý thức, chẳng hạn như khi mắc cỡ với câu hỏi về bệnh hoa liễu hay ngại vì câu trả lời sẽ bị chuyển đến cơ quan thuế. Cuối cùng, sai số do dụng cụ là việc cân chỉnh sai, thuốc thử không tinh khiết, thuốc thử pha loãng hay trộn không đúng cách hay do kiểm định chẩn đoán sai.

Độ nhạy cảm và độ đặc hiệu

Khả năng một thao tác phân loại đúng các cá nhân vào một trong hai nhóm, chẳng hạn như bệnh và không bệnh, tiếp xúc hay không tiếp xúc, dương tính hay âm tính, có nguy cơ cao hay không được đánh giá bởi 2 thông số độ nhạy cảm và độ đặc hiệu. Độ nhạy cảm (sensitivity) là tỉ lệ thực sự dương tính được chẩn đoán đúng và bằng 1 trừ cho tỉ suất âm tính giả (false negative rate). Độ đặc hiệu (specificity) là tỉ lệ thực sự âm tính được chẩn đoán đúng và bằng 1 trừ cho tỉ suất dương tính giả (false positive rate). Thí dụ, bảng 22.1 trình bày kết quả của một nghiên cứu dẫn đường để đánh giá khả năng bố mẹ nhớ lại đúng tình trạng tiêm chủng của trẻ, khi so với các hồ sơ sức khỏe chính thức. Trong 6 trẻ đã được tiêm chủng BCG, hầu hết, 55 trẻ được xác định đúng là đã tiêm chủng bởi cha mẹ của chúng, độ nhạy cảm tính ra là 55/60 hay 91,67%. Ngược lại 15 trong 40 trẻ không có hồ sơ về tiêm chủng BCG được bố mẹ cho là đã tiêm chủng, độ đặc hiệu là 25/40 hay 62,5%.

Bảng 22.1 So sánh sự nhớ lại của cha mẹ về tình trạng tiêm chủng của con với tình trạng tiêm chủng ghi trong hồ sơ sức khỏe chính thức.

Tiêm chủng BCG theo hồ sơ sức khỏe chính thức	Tiêm chủng BCG theo bố mẹ khai		
	Có	Không	Tổng số
Có	55	5	60
Không	15	25	40
Tổng số	70	30	100

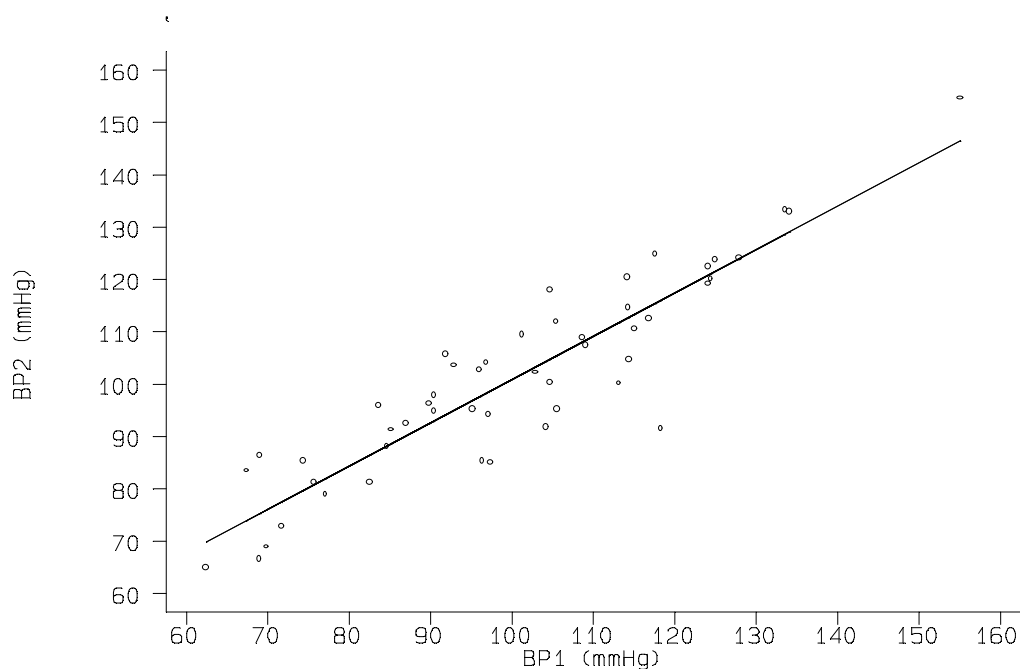
Số đo độ nhạy cảm và đặc hiệu rất quan trọng trong đánh giá một thử nghiệm sàng lọc (screening test). Cần lưu ý có mối liên hệ ngược giữa hai số đo này, xiết chặt (hay nói lỏng) tiêu chuẩn để cải thiện số đo này sẽ giảm độ lớn của số đo kia. Đường phân biệt chúng phụ thuộc vào bản chất của cuộc nghiên cứu. Thí dụ, trong thiết kế một cuộc nghiên cứu để thử nghiệm một loại vaccine bệnh phong mới, lúc đầu cần loại bỏ tất cả các bệnh nhân phong. Do đó người ta mong muốn một thử nghiệm có tỉ suất phát hiện dương tính thành công

cao, hay nói cách khác có độ nhạy cảm cao. Người ta không quan tâm nhiều đến độ đặc hiệu, bởi vì nó không ảnh hưởng gì nếu một trường hợp thực sự âm tính bị xác định lầm là dương tính và loại bỏ. Ngược lại, khi phát hiện các trường hợp mắc bệnh trong giai đoạn theo dõi sau khi tiêm chủng, người ta muốn một thử nghiệm có độ nhạy cảm cao, bởi vì điều quan trọng là phải tin chắc rằng tất cả các trường hợp dương tính được phát hiện là đúng, và nếu ít quan trọng hơn nếu chúng bị bỏ qua.

Hồi quy về trung bình

Hồi quy về trung bình nói về một hiện tượng đầu tiên được quan sát bởi Galton khi ông ta đầu tiên lưu ý rằng chiều cao của một số con trai có khuynh hướng dần tiến trung bình so với chiều cao của cha. Do đó người cha cao dễ có người con thấp hơn họ, trong khi đối với người ta thấp thì ngược lại. Không hiểu được điều này có thể dẫn tới lý giải sai trong hai trường hợp. Thứ nhất là hiệu số giữa hai lần đo lường có liên quan đến số đo lần thứ nhất. Thí nghiệm là trường hợp chỉ những cá nhân cực đoan trong phân phối được chọn điều trị và theo dõi. Điều này có thể lý giải bằng cách xét sự đo huyết áp lặp lại và đánh giá tác dụng của thuốc hạ áp lên huyết áp.

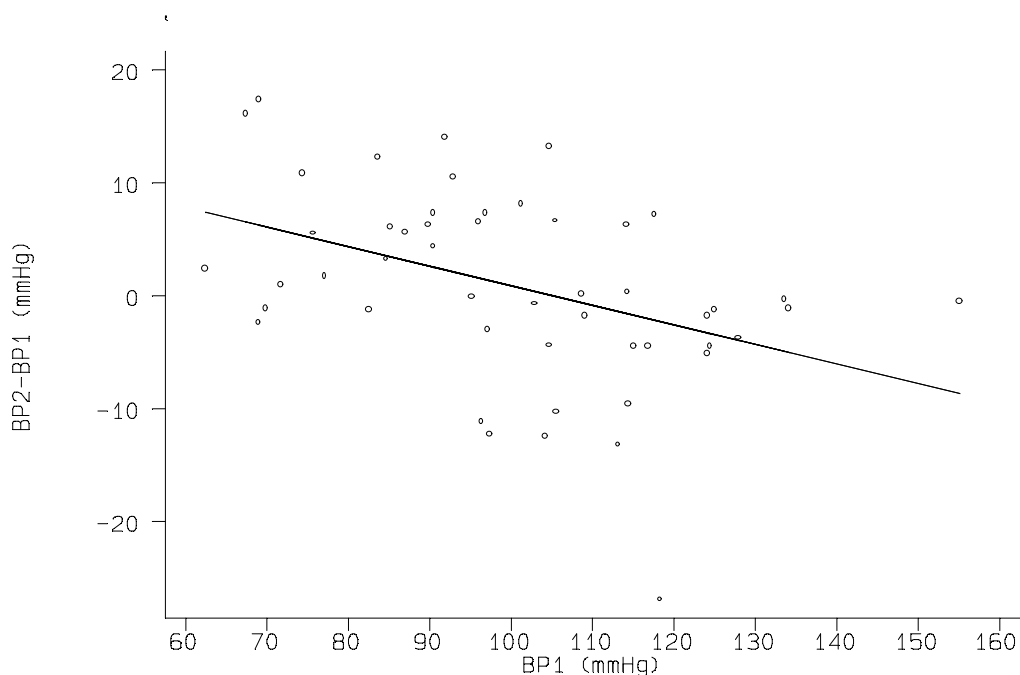
Hình 22.1 trình bày quan hệ giữa hai huyết áp tâm trương đọc cách nhau 6 tháng trên 50 người tình nguyện. Có thể thấy rằng ngoài sự biến thiên ngẫu nhiên có sự phù hợp giữa hai lần đo. Dù vậy bức tranh khi hiệu số giữa hai lần đọc được liên kết với trị số ban đầu sẽ khác hẳn (Hình 22.2). Có vẻ rằng có một khuynh độ đi xuống và những người có huyết áp ban đầu cao sẽ giảm huyết áp sau 6 tháng và ngược lại đối với người có huyết áp ban đầu thấp. Điều này hoàn toàn chỉ là sự liên hệ toán học giữa hiệu số (BP2 - BP1) và huyết áp lúc đầu BP1, bởi vì BP1 có mặt ở cả hai biến số.



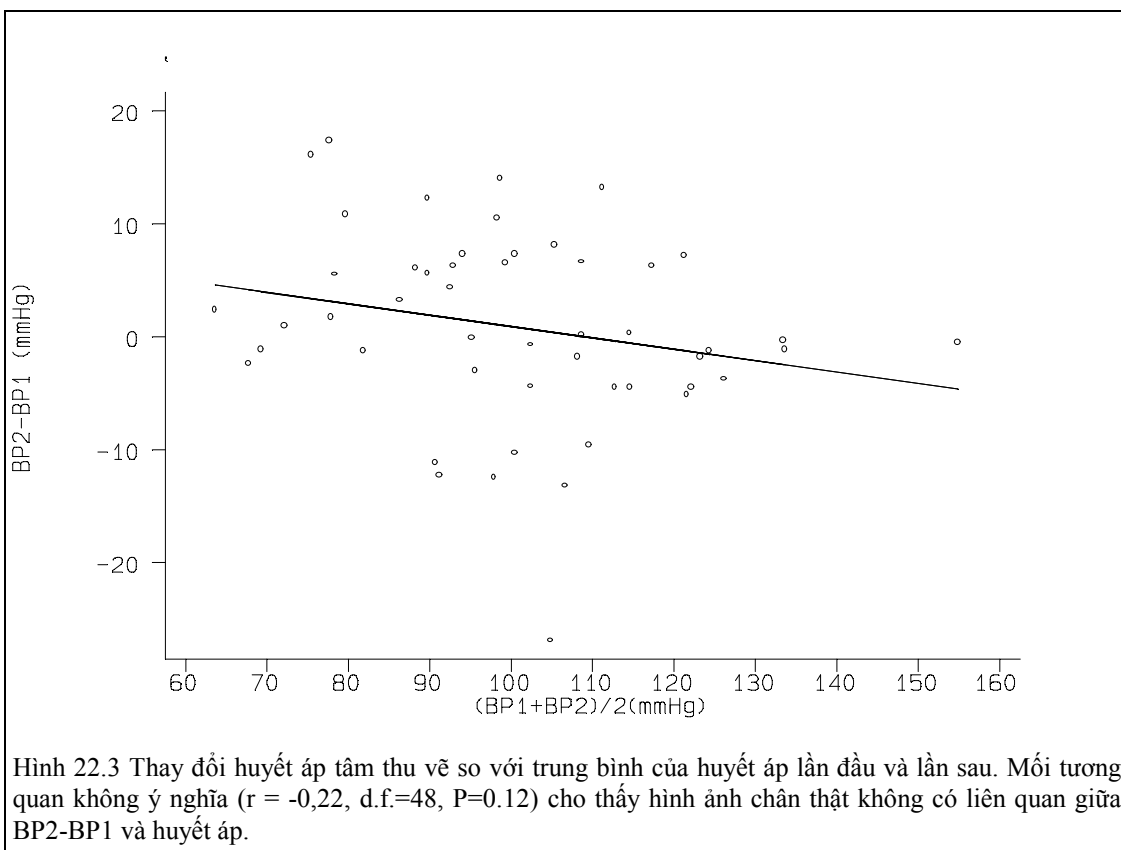
Hình 22.1 Quan hệ giữa hai huyết áp tâm trương cách nhau 6 tháng trên 50 người tình nguyện cho thấy sự thay đổi rất ít. Đường thẳng là quan hệ khi hai huyết áp trong hai trường hợp giống hệt nhau.

Đôi khi người ta muốn đánh giá có phải sự thay đổi trong điều trị có liên quan đến giá trị ban đầu. Chúng ta hãy xem xét ý nghĩa của hồi quy về trung bình đối với thử nghiệm lâm sàng của thuốc hạ huyết áp. Những thử nghiệm như vậy chỉ giới hạn ở những người có huyết áp tâm trương ban đầu cao, chẳng hạn như trên 120 mmHg. Có thể thấy từ hình 22.2 rằng những người được theo dõi trong một thời gian và đo lường lại huyết áp sẽ có số trung bình giảm đi,

ngay cả khi không có một trị liệu nào. Do đó cần có một nhóm chứng có huyết áp tương tự và để đánh giá mức giảm huyết áp trong nhóm điều trị bằng cách so sánh với mức giảm quan sát được trong nhóm chứng.



Hình 22.2 Thay đổi trong huyết áp tâm trương vẽ với giá trị ban đầu. Tương quan âm giả tạo ($r=-0,42$, $d.f.=48$, $P<0,01$). Đường thẳng là đường thẳng hồi quy tương ứng với sự liên hệ này.



Hình 22.3 Thay đổi huyết áp tâm thu vẽ so với trung bình của huyết áp lần đầu và lần sau. Mối tương quan không ý nghĩa ($r = -0,22$, $d.f.=48$, $P=0.12$) cho thấy hình ảnh chân thật không có liên quan giữa BP2-BP1 và huyết áp.

Để đánh giá xem tác dụng của thuốc có liên quan đến mức huyết áp hay không, có thể sử dụng phương pháp được đề nghị bởi Oldham (1962). Hiệu số $(BP1-BP2)$ được vẽ đối với trung bình của giá trị huyết áp lúc đầu và lúc sau $(BP1+BP2)/2$ chứ không phải với giá trị ban đầu. Tương quan âm tính của những biến số này có nghĩa là có khuynh hướng người có huyết cao giảm nhiều hơn những người có huyết áp thấp và tương quan dương tính có nghĩa ngược lại. Hình 22.3 trình bày đồ thị của $(BP1 - BP2)$ với $(BP1+BP2)/2$ cho số liệu của hình 22.1. Sự tương quan không có ý nghĩa có nghĩa rằng phương pháp Oldham khắc phục sự liên hệ giả tạo giữa giảm huyết áp và huyết áp ban đầu gây ra do sự hồi quy về trung bình.

PHƯƠNG PHÁP LẤY MẪU

Giới thiệu

Vai trò của điều trị cắt ngang và cắt dọc đã được phác họa trong Chương 21. Chúng ta sẽ mô tả các lược đồ lấy mẫu (sampling schemes) có thể được dùng. Đôi khi một cuộc tổng điều tra (census) dân số được tiến hành nhưng thường người ta chỉ nghiên cứu một mẫu (sample). Lý do chính là tiết kiệm thời gian và tiền bạc. Bất lợi chính là ước lượng dựa trên mẫu có thể bị sai số lấy mẫu, như trình bày trong Chương 1 và Chương 3, và có thể bị sai lệch nếu mẫu không đại diện cho dân số. Việc chọn lựa cỡ mẫu để nghiên cứu để cho sai số trong một giới hạn nhất định được mô tả ở Chương 26. Chọn lựa ngẫu nhiên mẫu được đề nghị để tránh các sai lệch chọn lựa, dù là có ý hay không, và các phương pháp thích hợp sẽ được phác thảo ngắn gọn. Dù vậy chỉ ngẫu nhiên hóa không đủ để khắc phục các sai lệch khả dĩ, như đã trình bày trong Chương 22. Thí dụ, sự không đáp ứng của một số cá nhân có thể gây sai lệch, bởi vì người không đáp ứng có thể khác với người đáp ứng về các khía cạnh cần nghiên cứu.

Chọn mẫu ngẫu nhiên đơn

Trong trong mẫu ngẫu nhiên đơn (simple random sampling), bước đầu tiên là rút ra một danh sách các cá nhân trong dân số và đánh số chúng. Danh sách này được gọi là khung mẫu (sampling frame). Sau đó chọn số người cần thiết; mỗi người có một cơ hội bằng nhau được chọn. Điều này có thể được chọn bằng cách đánh dấu một tấm thẻ cho mỗi cá nhân, xáo lên và chọn một số thẻ. Một phương pháp tiện lợi hơn là dùng bảng số ngẫu nhiên (table of random numbers) như trong bảng A10. Đầu tiên nhìn vào tổng số trong dân số và xem nó có bao nhiêu chữ số (digit). Sau đó chọn các số ngẫu nhiên có bấy nhiêu chữ số. Thí dụ, nếu dân số có 2432 phần tử thì cần 4 chữ số. Cuối cùng chọn một điểm khởi phát tùy ý trong bảng dự trữ (không nên luôn luôn khởi phát ở cùng một chỗ) và đọc các số lên, di chuyển xuống theo cột hay ngang theo hàng cho đến khi đã chọn được số lượng cần thiết các đối tượng khác nhau. Nếu số được chọn lớn hơn tổng số dân hay nếu bằng zero, bỏ qua số đó. Lưu ý rằng sự sắp xếp các số theo cặp trong bảng chỉ có mục đích trình bày mà thôi. Có thể dùng số ngẫu nhiên do máy tính tạo ra. Dù vậy không nên dùng các số ngẫu nhiên do các máy tính cầm tay tạo ra bởi vì chúng không có phân phối ngẫu nhiên thực sự.

Thí dụ 23.1

Thử chọn một mẫu gồm 10 nhà từ một làng có 268 nhà. Bảng 23.1 trình bày một phần của Bảng A10 để minh họa cách làm. Tổng số nhà là 268 do đó cần 3 chữ số. Chúng được tạo bằng cách dùng một cặp số và một chữ số đầu tiên của cặp kế bên. Bắt đầu từ góc trên bên trái và đọc theo cột xuống, số đầu tiên là 173. Số tiếp theo là 770 quá lớn và bỏ qua. Số tiếp theo trong phạm vi cần thiết là 074 hay 74. Điều này được tiếp tục cho đến khi chọn được 10 số khác nhau. Chúng được gạch dưới trong bảng. Do đó các nhà được nghiên cứu là 5, 74, 79, 83, 138, 166, 173, 201, 242, và 259.

Bảng 23.1 Một phần của bảng số ngẫu nhiên từ Bảng A10, minh họa sự lựa chọn 10 số giữa 1 và 268.

17	37	93	23	78	87	35	20	96	43
77	04	74	47	67	21	76	33	50	25
98	10	50	71	75	12	86	73	58	07
52	42	07	44	€	38	15	51	00	13
49	17	46	09	62	90	52	84	77	27
79	83	86	19	62	06	76	50	03	10
83	11	46	32	24	20	14	85	88	45
07	45	32	14	08	32	98	94	07	72
00	56	76	31	38	80	22	02	53	53
42	34	07	96	88	54	42	06	87	98
13	89	51	03	74	17	76	37	13	04
97	12	25	93	47	70	33	24	03	54
16	64	36	16	00	04	43	18	66	79
45	59	34	68	49	12	72	07	34	45
20	15	37	00	49	52	85	66	60	44

(Số 074 đã được chọn

Chọn mẫu hệ thống

Để tiện lợi, chọn từ một khung mẫu được tiến hành hệ thống thay vì ngẫu nhiên bằng cách chọn các cá nhân ở một khoảng cách nhất định trong danh sách, điểm bắt đầu được chọn ngẫu nhiên. Tỉ dụ để chọn mẫu 5% hay 1 phần 20 trong dân số, điểm khởi phát là được chọn ngẫu nhiên từ 1 đến 20, và sau đó cứ mỗi 20 người trong danh sách chọn một. Giả sử số ngẫu nhiên được chọn là 13 thì mẫu sẽ gồm các đối tượng 13, 33, 53, 73, 93 v.v.

Lấy mẫu hệ thống tương đương với lấy mẫu ngẫu nhiên với điều kiện không có một mô thức sắp các đối tượng trong mẫu. Dù vậy người ta thích lấy mẫu ngẫu nhiên bởi vì không cần giả thiết về cấu trúc dân số.

Các lược đồ lấy mẫu phức tạp hơn

Có một số tình huống trong đó lấy mẫu ngẫu nhiên đơn không thích hợp và cần các lược đồ lấy mẫu phức tạp hơn. Những tình huống đó phát sinh khi (i) không có khung mẫu và xây dựng khung mẫu rất tốn kém hay không thể được, (ii) dân số trải ra trong một khu vực rộng, chi phí và thời gian di chuyển để đi khắp khu vực không cho phép, (iii) dân số có những nhóm hoàn toàn khác biệt.

Có ba loại lược đồ lấy mẫu phức tạp cơ bản là lấy mẫu phân tầng (stratified), **hiệu bậc (multi-stage)** và **cụm (cluster)**. Chúng có thể dùng riêng hay kết hợp. Nói chung mục tiêu cần phải áp dụng cho những lược đồ này là làm sao cho mỗi cá nhân một cơ hội được chọn bằng nhau, đó là dùng chúng với phương pháp chọn **lựa xác suất bằng nhau (equal probability selection method - epsem)**.

Một điểm cần lưu ý là chỉ đối với các lược đồ epsem các giá trị của dân số chỉ được ước lượng trực tiếp bởi giá trị lấy mẫu chung (overall sample value). Nếu sử dụng các phương pháp xác suất không bằng nhau, phương pháp ước lượng giá trị dân số sẽ khác (xem Thí dụ 23.2). Trong cả hai trường hợp, việc tính toán sai số chuẩn sẽ phụ thuộc vào lược đồ lấy mẫu. Để thảo luận đầy đủ về các vấn đề thực tiễn có liên quan, người đọc cần tham khảo Lutz (1986) và Moser và Kalton (1971) và các chi tiết lý thuyết ở Cochran (1977) và Stuart (1984).

Lấy mẫu phân tầng

Lấy mẫu phân tầng được dùng khi dân số bao gồm các nhóm khác biệt, hay tầng (strata), khác nhau về các đặc tính nghiên cứu và bản thân chúng cũng cần quan tâm. Những thí dụ thường gặp là các nhóm tuổi giới tính hay các vùng khác nhau trong quốc gia. Một mẫu ngẫu nhiên đơn được rút ra từ mỗi tầng để đảm bảo rằng chúng đủ đại diện. Ước lượng chung chính xác hơn dựa vào ngẫu nhiên đơn không tính đến cấu trúc các nhóm nhỏ trong dân số. Chiến lược thường dùng là chọn các cá nhân trong mỗi tầng với tỉ lệ như nhau (lọc đồ epsem), nghĩa là dùng cùng một **phân số lấy mẫu (sampling fraction) cho mỗi tầng. Dù vậy, đôi khi cần phải** thay đổi để cho cỡ mẫu trong mỗi tầng không quá nhỏ.

Bảng 23.2 Kết quả một mẫu phân tầng được tiến hành để ước lượng tỉ suất mắc toàn bộ của một bệnh trong một quốc gia có ba vùng địa lí chính. Tỉ suất mắc toàn bộ chung được tính bằng cách cộng số các người bị bệnh ước lượng được trong mỗi vùng và chia cho tổng số dân.

Khu vực	dân số	cỡ mẫu	số bị bệnh	tỉ suất mắc toàn bộ	tổng số bị bệnh ước lượng
Đồng bằng ven biển	150000	200	120	0,6	900 000
Vùng núi	150000	50	5	0,1	15 000
Bán hoang mạc	300000	50	15	0,3	90 000
Tổng số	1950000	300	140	*0,52	1 005 000

Thí dụ 23.2

Người ta muốn ước lượng tỉ suất mắc toàn bộ của bệnh trong một quốc gia với ba vùng địa lí chính, vùng đồng bằng ven biển, vùng núi và vùng bán hoang mạc. Bởi vì dân số phân phối đồng đều trong quốc gia, và bởi vì người ta nghĩ rằng địa lí có thể ảnh hưởng đến tỉ suất mắc toàn bộ của bệnh, người ta chọn một mẫu phân tầng. Bảng 23.2 trình bày kết quả thu được với tỉ suất mắc toàn bộ trong mỗi vùng.

Tỉ suất mắc toàn bộ chung được tính bằng cách ước lượng số người bị bệnh trong mỗi vùng. Thí dụ trong vùng đồng bằng ven biển tỉ suất mắc toàn bộ của mẫu là $120/200$ hay $0,6$. Áp dụng số này cho tổng số dân số trong vùng đồng bằng ven biển cho số ước lượng $0,6 \times 1500000 = 900000$ người bị bệnh. Số người bị bệnh của vùng núi và bán hoang mạc được tính tương tự là 15000 và 90000 . Tổng số người bị bệnh trong toàn quốc gia là 1050000 . Kích thước dân số là 1950000 cho nên tỉ suất mắc toàn bộ chung là $1050000/1950000 = 0,52$.

Lưu ý rằng số này không giống như tỉ suất mắc toàn bộ của mẫu là $140/300 = 0,47$. Hai kết quả chỉ giống nhau khi dùng phân số lấy mẫu giống nhau cho mỗi tầng (nhưng không xảy ra trong trường hợp này). Việc tính toán sai số chuẩn của tỉ suất mắc toàn bộ cho toàn dân số dựa trên sự kết hợp các sai số chuẩn của các tỉ suất mắc toàn bộ của mỗi vùng. Xem Moser và Kalton (1971) để biết thêm chi tiết.

Lấy mẫu nhiều bậc

Lấy mẫu nhiều bậc (multi-stage sampling) được tiến hành trong nhiều bậc dùng các cấu trúc đẳng cấp (hierarchical structure) của dân số. Thí dụ, lấy mẫu hai bậc (two stage sampling) có thể bao gồm lần thứ nhất lấy một mẫu ngẫu nhiên các trường hợp và sau đó lấy một mẫu ngẫu nhiên các trẻ em trong các trường được chọn. Các trường được gọi là đơn vị bậc một (first stage units) và trẻ em là đơn vị **bậc hai (second-stage units)**. **Ưu điểm là tài nguyên có thể được tập trung tại một số** địa điểm và không cần cơ cấu lấy mẫu cho toàn dân số. Cần danh sách các đơn vị bậc một nhưng chỉ cần danh sách các đơn vị bậc hai của các đơn vị bậc một đã được chọn. Khuyết điểm là ước lượng chung kém chính xác hơn khi dựa trên lấy mẫu ngẫu nhiên đơn có cùng một cỡ mẫu chung. Nó khác, để đạt được cùng độ chính xác như lấy mẫu ngẫu nhiên đơn cần một cỡ mẫu lớn hơn.

Lấy mẫu ở bậc hai gồm lấy các mẫu ngẫu nhiên đơn có cùng cỡ từ các đơn vị bậc một. Phương pháp lấy mẫu bậc một phụ thuộc vào chúng có chứa một số các đơn vị bậc hai như nhau hay không. Nếu có, có thể mẫu mẫu ngẫu nhiên đơn. Nếu chúng có cỡ khác nhau, có thể đạt được lược đồ eptem, bằng cách lấy mẫu xác suất tỉ lệ **với kích thước (probability proportional to size - p.p.s)**. **Thí dụ, nếu một trường** có nhiều gấp đôi học sinh so với trường kia nó có cơ hội được chọn gấp đôi. Lấy mẫu p.p.s. Được tiến hành hoàn lại mẫu (with replacement), có nghĩa là sau khi một đơn vị bậc một được chọn nó vẫn còn được rút chọn và có thể được chọn lần nữa. Khi một đơn vị bậc một được chọn hai lần, chọn mẫu đơn vị bậc hai nhiều gấp đôi. Tác dụng chung là cho mỗi đơn vị bậc hai trong dân số một cơ hội được chọn bằng nhau.

Có thể có lược đồ lấy mẫu có hơn hai mức lấy mẫu, thí dụ như chọn tỉnh, quận, đường phố và cuối cùng là nhà. Nó được gọi là lấy mẫu nhiều bậc (multi-stage sampling).

Lấy mẫu cụm

Nếu chi phí phụ trội không nhiều, nên điều tra tất cả các đơn vị bậc hai từ mỗi đơn vị bậc một được chọn trong lược đồ lấy mẫu hai bậc. Điều đó được gọi là lấy mẫu **cụm (cluster sampling)** và **đơn vị lấy mẫu bậc một được gọi là cụm (cluster)** trong trường hợp này. Có thể đạt được lược đồ xác suất bằng nhau bằng cách lấy mẫu ngẫu nhiên đơn các cụm bất kể rằng chúng có kích thước bằng nhau hay không.

Lấy mẫu cụm được dùng nếu có lợi ích được phân phát cho người tham gia và nếu không thích hợp hoặc không đạo đức nếu chỉ phân phát cho một số thành viên của đơn vị. Thí dụ, trong khi lấy mẫu trường để ước lượng tỉ suất mắc toàn bộ của bệnh khi muốn sử dụng một phương pháp điều trị hiệu quả cho tất cả những người bị bệnh, người ta sẽ muốn khám tất cả các học sinh trong các trường được chọn chứ không khám một mẫu trong đó.

Thí dụ 23.3

Lấy mẫu phân tầng được đề nghị trong thí dụ 23.2 để ước lượng tỉ suất mắc toàn bộ trong một quốc gia với 3 vùng chính có thể được cải tiến thành mẫu cộng đồng thứ nhất (thành phố, làng, ấp) và các hàn tong vùng, khám tất cả các thành viên trong nhà. Lược đồ sẽ là sự kết hợp giữa lấy mẫu phân tầng (khu vực) lấy mẫu hai bậc (cộng đồng và nhà) và lấy mẫu cụm (tất cả các thành viên trong nhà).

NGHIÊN CỨU ĐOÀN HỆ VÀ BỆNH CHỨNG

Giới thiệu

Ưu và khuyết điểm của phương pháp bệnh chứng so với phương pháp đoàn hệ (hay cắt dọc) đã được thảo luận ở chương 21. Một nghiên cứu đoàn hệ gồm theo dõi các cá nhân từ khi xác định tình trạng tiếp xúc đến sự xuất hiện bệnh sau đó. Ngược lại, nghiên cứu bệnh chứng bắt đầu với tình trạng bệnh tật của cá nhân và nhìn ngược lại thời gian để xác định tình trạng tiếp xúc trong quá khứ đối với yếu tố nguy cơ quan tâm. Vì lý do này các từ tiền cứu (prospective) và hồi cứu (retrospective) thường được dùng lẫn lộn với đoàn hệ và bệnh chứng. Dù vậy cách dùng từ này gây nhầm lẫn và không nên dùng bởi vì như đã lưu ý trong Chương 21, đôi khi có thể tiến hành nghiên cứu đoàn hệ hoàn toàn hồi cứu trong các hồ sơ cũ.

Trong chương này chúng ta sẽ tập trung vào việc đo lường mối liên hệ dùng để đánh giá sức mạnh của quan hệ giữa yếu tố nguy cơ và sự xuất hiện bệnh sau đó. Nghiên cứu bệnh chứng được hiểu dễ dàng bằng cách hiểu rằng đằng sau mỗi nghiên cứu bệnh chứng có một nghiên cứu đoàn hệ tương đương. Do đó chúng ta xem nghiên cứu đoàn hệ trước.

Bảng 24.1 Số liệu từ một cuộc nghiên cứu đoàn hệ để điều tra mối liên hệ giữa hút thuốc lá và ung thư phổi. Có minh họa tính toán nguy cơ tương đối và quy trách.

	ung thư phổi	không ung thư phổi	tổng số	Tỉ suất mới mắc
Người hút thuốc	39	29961	300000	1,3/1000/năm
Không hút thuốc	6	59994	60000	0,1/1000/năm
Tổng số	45	89955	90000	
$RR = \frac{1,30}{1,10} = 13,0$ $AR = 1,3 - 0,1 = 1,2/1000/năm$ $AR\% = \frac{1,20}{1,30} = 0,923 = 92,3\%$				

Nghiên cứu đoàn hệ

Trong nghiên cứu đoàn hệ, một mẫu các cá nhân, một số tiếp xúc với yếu tố nguy cơ quan tâm và một số không, được theo dõi theo thời gian và tỉ suất mắc bệnh của hai nhóm được so sánh với nhau. Thí dụ, Bảng 24.1 cho thấy một số liệu từ một nghiên cứu đoàn hệ để điều tra mối liên hệ giữa hút thuốc lá và ung thư phổi. Ba mươi ngàn người hút thuốc và 60000 ngàn người không hút thuốc được theo dõi trong một năm, trong thời gian đó 39 người hút thuốc và 6 người không hút thuốc bị ung thư phổi, tính ra nguy cơ mới mắc tương ứng là 1,3 và 0,1 cho mỗi 1000 người mỗi năm. Do đó nguy cơ mới mắc ung thư phổi cao hơn đáng kể trong những người hút thuốc so với người không hút thuốc. Trên thực tế chúng ta thấy 13 lần cao hơn. Tỉ số này được gọi là nguy cơ tương đối (relative risk) và tổng kết được sức mạnh của mối liên hệ giữa yếu tố và bệnh tật.

Nguy cơ tương đối

$$\text{Nguy cơ tương đối (RR)} = \frac{\text{Tỉ suất mới mắc trong nhóm tiếp xúc}}{\text{Tỉ suất mới mắc trong nhóm không tiếp xúc}}$$

Nguy cơ tương đối (Relative risk - RR) bằng 1 xảy ra khi tỉ suất mới mắc như nhau trong hai nhóm và tương đương với không có mối liên hệ giữa yếu tố nguy cơ và bệnh. Nguy cơ tương đối lớn hơn 1 khi nguy cơ mắc bệnh cao ở nhóm tiếp xúc với yếu tố hơn nhóm không tiếp xúc với yếu tố, như trong thí dụ trên với tiếp xúc là hút thuốc lá. Nguy cơ tương đối nhỏ hơn một khi nguy cơ trong nhóm tiếp xúc thấp hơn, gợi ý rằng yếu tố có thể là bảo vệ. Một thí dụ là nguy cơ tiêu chảy giảm được quan sát trong những trẻ bú sữa mẹ so với những trẻ không bú.

Nguy cơ tương đối càng xa một bao nhiêu thì liên hệ càng mạnh bấy nhiêu. Ý nghĩa thống kê có thể được kiểm định bằng cách dùng kiểm định χ^2 2×2 được mô tả ở Chương 13.

Khoảng tin cậy của RR

Việc tính toán khoảng tin cậy khá phức tạp. Công thức ở đây là của Miettinen. Nó là xấp xỉ dựa trên kiểm định có nghĩa là nó được tính dựa trên giá trị χ^2 (trên thực tế là căn của χ^2) tìm được trong kiểm định ý nghĩa chứ không phải là sai số chuẩn.

$$95\% \text{ c.i.} = RR^{(1 \pm 1,96/\chi)}$$

Có thể tính được khoảng tin cậy phần trăm khác bằng cách thay 1,96 với điểm phần trăm thích hợp trong phân phối bình thường.

Thí dụ, xem số liệu được trình bày trong bảng 24.1

$$RR = 13,0, \chi^2 = 52,25, P < 0,001$$

Do đó

$$\chi = (52,25 = 7,43 \text{ và } 1,96/\chi = 1,96/7,43 = 0,26$$

Nên

$$RR^{(1+1,96/\chi)} = 13,0^{1+0,26} = 13,0^{1,26} = 25,3$$

Và

$$RR^{(1-1,96/\chi)} = 13,0^{1-0,26} = 13,0^{0,74} = 6,7$$

Khoảng tin cậy 95% của nguy cơ tương đối do đó là 6,7 đến 25,3

Nguy cơ quy trách

Nguy cơ tương đối đánh giá, chẳng hạn như một người hút thuốc lá có thể bị ung thư phổi gấp bao nhiêu lần người không hút thuốc, nhưng không nói gì đến độ lớn của nguy cơ vượt quá theo số tuyệt đối. Số này được đo lường bởi nguy cơ quy trách (**attributable risk - AR**).

$$\text{Nguy cơ quy trách (AR)} = \frac{\text{tỉ suất mắc bệnh trong nhóm tiếp xúc} - \text{tỉ suất mắc bệnh trong nhóm không tiếp xúc}}$$

Nguy cơ quy trách đôi khi được tính bằng tỉ lệ so với tổng số tỉ suất mới mắc trong nhóm tiếp xúc và được gọi là phần trăm nguy cơ quy trách (attributable risk percent) hay nguy cơ quy trách tỉ lệ (proportional attributable risk), tỉ lệ quy trách (tiếp xúc), phân số quy trách (tiếp xúc) và phân số căn nguyên (tiếp xúc).

$$\begin{aligned} \text{Nguy cơ quy trách tỉ lệ (AR\%)} &= \frac{\text{tỉ suất mắc bệnh trong nhóm tiếp xúc} - \text{tỉ suất mắc bệnh trong nhóm không tiếp xúc}}{\text{tỉ suất mắc bệnh trong nhóm tiếp xúc}} \\ &= \frac{RR - 1}{RR} \end{aligned}$$

Trong thí dụ này, nguy cơ quy trách của ung thư phổi do hút thuốc lá là $13,0 - 1,0 = 12$ trường hợp cho mỗi 1000 mỗi năm, với hút thuốc lá chịu trách nhiệm cho 92,3% ($12/13$) tất cả các trường hợp ung thư phổi trong những người hút thuốc.

Bảng 24.2 Nguy cơ tương đối và quy trách của tử vong do các nguyên nhân khác nhau, 1951 đến 1961, liên hệ với hút thuốc lá nặng ở các nam bác sĩ người Anh. Số liệu từ Doll & Hill (1964), British medical journal 1, 1399-1410, được trình bày bởi MacMahon & Pugh (1970) Epidemiology, Principles and Methods. Little, Brown & Co., Boston (đã xin phép)

Nguyên nhân chết	tỉ suất tử vong đặc hiệu tuổi (cho mỗi 1000 người-năm)		RR	AR
	Không hút thuốc lá	Hút thuốc nặng		
Ung thư phổi	0,07	2,27	32,4	2,20
Các ung thư khác	1,91	2,59	1,4	0,68
Viêm phế quản mãn	0,05	1,06	21,2	1,01
Bệnh tim mạch	7,32	9,93	1,4	2,61
Các nguyên nhân	12,06	19,67	1,6	7,61

Bảng 24.2 trình bày nguy cơ tương đối và nguy cơ quy trách của tử vong từ các nguyên nhân khác nhau liên hệ đến hút thuốc lá nặng. Mỗi liên hệ được minh họa rõ nhất đối với ung thư phổi và viêm phế quản mãn tính, với nguy cơ tương đối tương ứng là 32,1 và 21,2. Dù vậy, nếu sự liên hệ với bệnh tim mạch, dù là không mạnh như thế, cũng được xem có tính nhân quả, loại bỏ hút thuốc lá sẽ cứu được nhiều bệnh nhân chết vì tim mạch hơn là vì ung thư phổi, 2,61 so với 2,20 cho mỗi 1000 người hút thuốc lá mỗi năm. Lưu ý rằng tỉ suất tử vong được chuẩn hóa để xét đến sự khác biệt do phân phối giữa người hút thuốc và người không hút thuốc, và sự gia tăng tỉ suất chết theo tuổi.

Nói tóm lại, nguy cơ tương đối là đo lường tốt nhất giữa sức mạnh của sự liên hệ giữa yếu tố nguy cơ và bệnh tật. Nó càng lớn thì càng có thể rằng mối liên quan là nhân quả. Mặt khác, nguy cơ quy trách cho ý niệm tốt hơn về nguy cơ thặng dư của bệnh do một cá nhân gánh chịu do sự tiếp xúc. Dù vậy, tổng số tác động của yếu tố nguy cơ trong dân số sẽ phụ thuộc sự tiếp xúc phổ biến như thế nào. Tính theo dân số, sự tiếp xúc hiếm với nguy cơ tương đối cao có thể ít quan trọng hơn một tiếp xúc rất phổ biến với nguy cơ tương đối thấp hơn. Điều này có thể được đánh giá bằng tỷ suất mới mắc chung trong dân số so với tỉ suất mới mắc trong nhóm không tiếp xúc. Số này gọi là nguy cơ quy trách dân số (population attributable risk).

$$\text{Nguy cơ quy trách dân số} = \frac{\text{tỉ suất mới mắc trong dân số chung} - \text{tỉ suất mới mắc trong nhóm không tiếp xúc}}{\text{tỉ suất mới mắc trong dân số chung}}$$

Nó có thể được tính bằng tỉ lệ của tỉ suất mới mắc chung. Số này được gọi là **phân số quy trách dân số (population attributable fraction)**, tên gọi khác là **phân số căn nguyên** (dân số) hay nguy cơ quy trách tỉ lệ dân số (population proportional attributable risk).

$$\begin{aligned} \text{Phân số quy trách dân số} &= \frac{\text{tỉ suất mới mắc trong dân số chung} - \text{tỉ suất mới mắc trong nhóm không tiếp xúc}}{\text{tỉ suất mới mắc trong dân số chung}} \\ &= \frac{\text{tỉ suất mới mắc trong dân số chung} (RR - 1)}{1 + \text{tỉ suất mới mắc trong dân số chung} (RR - 1)} \end{aligned}$$

Tỉ suất mới mắc, nguy cơ mới mắc và tỉ số số chênh

Nhớ lại từ Chương 15 rằng số mới mắc có thể là nguy cơ hay một tỉ suất. Việc tính toán nguy cơ mới mắc dựa trên dân số nguy cơ ở lúc đầu nghiên cứu, trong khi tỉ suất dựa vào tổng số người năm nguy cơ trong nghiên cứu và phản ánh sự thay đổi của dân số nguy cơ. Điều này được minh họa trong hình 24.1 phản ánh sự tích lũy dần dần các trường hợp bệnh từ những người không bị bệnh ở trong dân số tiếp xúc và không tiếp xúc.

Một cách khác để đo lường mức độ mắc là số chênh (odds) của bệnh so với không bệnh. Nó bằng tổng số các trường hợp bệnh chia cho tổng số người nguy cơ ở cuối nghiên cứu. Dùng kí hiệu trong bảng 24.3, số chênh trong nhóm tiếp xúc là a/b và trong nhóm không tiếp xúc là c/d . Tỉ số này được gọi là tỉ số số chênh (odds ratio). Nó bằng;

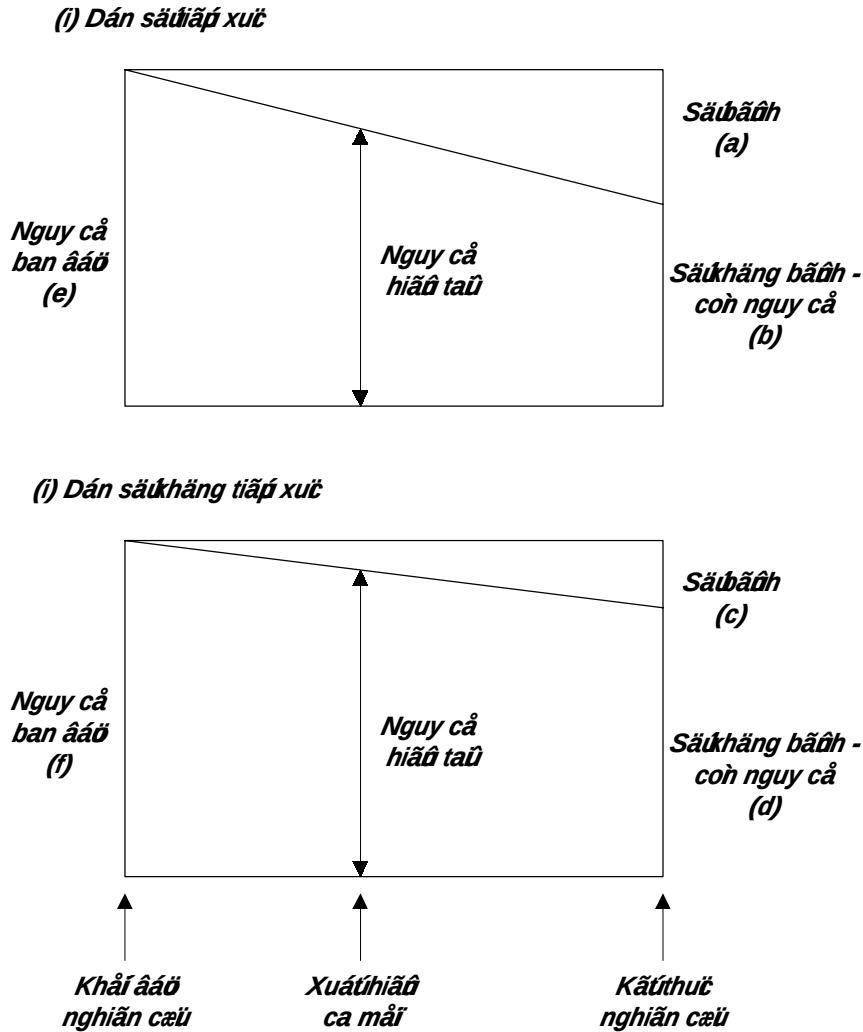
$$\frac{a/b}{c/d} = \frac{ad}{bc}$$

Là tích chéo của bảng. Tỉ số số chênh được ước tính từ bảng 24.1 bằng $(39 \times 59994) / (6 \times 29961) = 13,0$

Tỉ số số chênh (OR) = ad/bc

Tỉ số số chênh cũng có thể được xem là tỉ số của số chênh tiếp xúc và không tiếp xúc trong nhóm bệnh (a/c) so với nhóm không bệnh (b/d). Vì lí do này nó đóng vai trò quan trọng trong nghiên cứu bệnh- chứng (xem ở dưới).

Trong các bệnh hiếm, đó là bệnh mà tỉ lệ các trường hợp bệnh trong dân số rất thấp, đường trong hình 24,1 cho thấy sự tích lũy của các trường hợp bệnh sẽ hầu như nằm ngang và không thể phân biệt với đường ở trên thể hiện dân số nghiên cứu. Trong trường hợp này nguy cơ mới mắc, tỉ suất mới mắc và tỉ số số chênh bằng nhau.



Hình 24.1 Một trình bày hình ảnh một nghiên cứu đoàn hệ. Có 3 số đo của mối mắc tương đối là:

$$\begin{aligned} \text{Nguy cả tắắắ áắắ} &= \frac{a/e}{c/f} \\ \text{(Sắắ duắắ nguy cả mắắ mắắ)} & \\ \text{Nguy cả tắắắ áắắ} &= \frac{a/(\text{pyar tiếấ xưế})}{b/(\text{pyar kắắắ tiếấ xưế})} \\ \text{(Sắắ duắắ tề sắắắ mắắ mắắ)} & \\ \text{Tề sắắắắắắ} &= \frac{a/b}{c/d} = \frac{ad}{bc} \end{aligned}$$

Dù vậy, đối với các bệnh phổ biến, ba số này khác nhau và sẽ cho 3 số đo mối liên hệ giữa yếu tố và bệnh khác nhau, nguy cơ tương đối dùng nguy cơ mới mắc, nguy cơ tương đối dùng tỉ suất mới mắc và tỉ số số chênh. Cách thông thường là chọn tỉ suất mới mắc. Dù vậy việc chọn nguy cơ mới mắc thích hợp hơn khi đánh giá tác dụng bảo vệ của sự tiếp xúc, thí dụ như vaccine mà người ta tin rằng nó bảo vệ đầy đủ với một số cá nhân nhưng không đối với các cá nhân khác, chứ không phải bảo vệ một phần cho tất cả (Smith et al. 1984). Tỉ số số chênh được dùng trong nghiên cứu bệnh chứng (xem ở dưới).

Bảng 24.3 Kí hiệu cho nghiên cứu đoàn hệ và nghiên cứu bệnh chứng không bắt cặp

	Bệnh	không bệnh	tổng số
Tiếp xúc	a	b	e
Không tiếp xúc	c	d	f
Tổng số	g	h	n

Nghiên cứu bệnh chứng

Trong nghiên cứu bệnh chứng (case-control study), mẫu được lấy theo tình trạng bệnh chứ không phải theo tình trạng tiếp xúc. Một nhóm người được xác định có bệnh, nhóm bệnh (cases), được so sánh với một nhóm người không có bệnh, nhóm chứng (control) về phương diện tiếp xúc trước đó đối với yếu tố quan tâm. Không thể trực tiếp thu được thông tin về bệnh mới mắc trong dân số tiếp xúc và không tiếp xúc, nhưng có thể thấy rằng tỉ số tích chéo (cross product ratio) trong bảng bệnh chứng ước lượng tỉ số số chênh. Đối với bệnh hiếm, nó bằng với nguy cơ tương đối.

Bảng 24.4 trình bày kết quả của một nghiên cứu bệnh chứng tìm hiểu mối liên quan giữa uống cà phê và ung thư tuyến tụy. Người ta thu được nguy cơ tương đối là 2,5, có nghĩa là phụ nữ uống cà phê dễ bị ung thư gấp 2,5 lần hơn những người không uống.

Biểu hiấ: $OR = ad/bc = RR$

Biểu phẩ: $OR = ad/bc, 0 < |RR| < |OR|$

Bảng 24.4 Kết quả nghiên cứu bệnh chứng tìm hiểu sự liên hệ giữa uống cà phê và ung thư tuyến tụy. Số liệu ở phụ nữ được trình bày. Với sự cho phép của MacMahon và cộng sự (1981) New England Journal of Medicine 304,630-3.

Uống cà phê	Bệnh	Chứng	Tổng số
Có	140	280	420
Không	11	56	67
Tổng số	151	336	487

$$OR = \frac{140 \times 56}{11 \times 280} = 2,5$$

Đối với bệnh phổ biến, ý nghĩa của tỉ số tích chéo phụ thuộc vào lược đồ lấy mẫu được dùng cho nhóm chứng (Smith et al. 1984). Người ta thường chọn nhóm chứng từ những người không bị bệnh ở cuối cuộc nghiên cứu; bất cứ người chứng được chọn trong quá trình nghiên cứu sau đó mắc bệnh được coi là bệnh chứ không phải là chứng. Trong trường hợp này có thể thấy rằng tỉ số tích chéo đại diện cho tỉ số số chênh, và độ lớn của tỉ số số chênh $|OR|$ lớn hơn độ lớn của nguy cơ tương đối.

Các biến gây nhiễu

Trong quy hoạch và phân tích các nghiên cứu bệnh chứng, điều quan trọng là đảm bảo không có biến số gây nhiễu có thể gây sai lệch trong nghiên cứu mối liên quan. Thí dụ, trong nghiên cứu uống cà phê và ung thư tuyến tụy, cần thiết đảm bảo rằng phân phối tuổi của 2 nhóm bệnh và chứng tương tự nhau, bởi vì nguy cơ ung thư tuyến tụy gia tăng với tuổi. Điều này có thể sắp xếp bằng cách hoặc là chọn số bệnh bằng số chứng trong mỗi nhóm tuổi (bắt cặp tầng - stratum matching) hay dùng một quy hoạch bắt cặp (xem ở dưới). Một cách khác, có thể xét đến biến số gây nhiễu trong khi phân tích, thí dụ như chia số liệu thành những nhóm tuổi khác nhau và tính tỉ số số chênh cho mỗi nhóm. Tính ước lượng chung bằng cách tính ad/n và bc/n trong mỗi bảng. Những ad/n này được cộng với nhau và chia cho tổng số các bc/n . Số

này được gọi là ước lượng Mantel-Haenszel của tỉ số chênh và kiểm định ý nghĩa thích hợp là kiểm định χ^2 tổng kết Mantel-Haenszel (xem Chương 14).

$$OR_{\text{Mantel-Haenszel}} = \frac{\sum ad / n}{\sum bd / n}$$

Một cách tiếp cận tinh vi hơn là dùng hồi quy logistic (xem Chương 14), và để ước lượng tỉ số chênh từ hệ số hồi quy. Hồi quy này thích hợp để kiểm soát đồng thời nhiều yếu tố gây nhiễu, trong trường hợp này số các nhóm nhỏ lớn và số các cá nhân trong mỗi nhóm lại rất nhỏ. Xem chi tiết ở Breslow và Day (1980) và Schlesselman (1982).

Quy hoạch bắt cặp

Trong quy hoạch bắt cặp (matched design), mỗi bệnh được bắt cặp với một hay nhiều chứng, được chọn sao cho các biến số gây nhiễu có cùng giá trị. Thí dụ, Bảng 24.5 trình bày số liệu từ một nghiên cứu tìm hiểu mối liên hệ giữa sử dụng thuốc ngừa thai và thuyên tắc mạch. Bệnh gồm 175 phụ nữ tuổi từ 15-44 được xuất viện từ 43 bệnh viện sau khi cơn thuyên tắc mạch đầu tiên. Một bệnh nhân nữ bị một bệnh khác (không liên quan đến việc sử dụng thuốc ngừa thai) được chọn lọc từ cùng một bệnh viện để làm chứng cho mỗi bệnh. Người chứng được chọn có cùng nơi cư ngụ, cùng thời gian nhập viện, chủng tộc, tuổi tác, tình trạng hôn nhân, tình trạng thai sản và tình hình thu nhập y như người bệnh. Người tham gia được hỏi về tiền căn dùng thuốc ngừa thai và đặc biệt là họ có dùng thuốc ngừa thai vào tháng trước nhập viện hay không.

Bảng 24.5 Kết quả nghiên cứu bệnh chứng có bắt cặp số liệu để điều tra mối quan hệ giữa dùng thuốc ngừa thai (OC) và thuyên tắc mạch. Được phép từ Sartwell et al. (1969) American Journal of Epidemiology 90, 365-80

		Chứng		
		Dùng OC	Không dùng OC	Tổng số
Bệnh	Dùng OC	10	57	67
	Không dùng OC	13	95	108
	Tổng số	23	152	175
		$OR = \frac{57}{13} = 4,4$		

Cặp đôi bệnh và chứng được bảo tồn trong quá trình phân tích bằng việc lập bảng việc sử dụng thuốc ngừa thai trong bệnh so với sử dụng thuốc ngừa thai trong nhóm chứng. Có 10 cặp bệnh chứng trong đó cả bệnh và chứng đều dùng thuốc ngừa thai và 95 cặp đều không dùng thuốc ngừa thai. 105 cặp này không cho thông tin nào về sự liên hệ. Thông tin hoàn toàn chứa đựng trong 70 cặp mà bệnh, chứng khác nhau. Có 57 cặp bệnh chứng trong đó chỉ có bệnh dùng thuốc ngừa thai trong tháng trước so với 13 cặp trong đó chỉ có chứng dùng thuốc ngừa thai trong tháng trước. Tỉ số chênh được đo bằng tỉ số những cặp không phù hợp (discordant pairs) và bằng 4,4 (= 57/13). Kiểm định ý nghĩa thích hợp là kiểm định χ^2 McNemar để so sánh tỉ lệ cặp đôi (xem Chương 14), cho $\chi^2 = 26,4$, $P < 0,001$.

OR = tỉ số chênh cặp kháng phù hợp

$$= \frac{\text{Số cặp trong đó bệnh tiếp xúc, chứng không}}{\text{Số cặp trong đó chứng tiếp xúc, bệnh không}}$$

Nếu chọn nhiều chứng cho một bệnh, tỉ số chênh được ước lượng bằng cách áp dụng thao tác Mantel-Haenszel. Cuối cùng, việc phân tích một số yếu tố nguy cơ, hay cần điều chỉnh các biến số gây nhiễu khác ngoài các biến số gây nhiễu đã được bắt cặp trong quy hoạch, cần một dạng hồi quy đặc biệt là hồi quy logistic có điều kiện (conditional logistic regression).

Những phương pháp khác nhau được trình bày trong bảng 24.6. Thảo luận đầy đủ trong những tình huống phức tạp hơn, tham khảo Breslow và Day (1980) hay Schlesselman (1982).

Bảng 24.6 Phân tích nghiên cứu bệnh chứng - tổng kết các phương pháp

Lược đồ lấy mẫu cho chứng	Một yếu tố nguy cơ	Nhiều yếu tố nguy cơ/điều chỉnh cho các biến số gây nhiễu
<i>(a) một chứng cho một bệnh</i>		
Ngẫu nhiên	Bảng 2×2 đơn trình bày yếu tố nguy cơ × bệnh/chứng Kiểm định χ^2 chuẩn OR=tỉ số tích chéo, ad/bc	hồi quy logistic hay phân tích phân tầng
Bất cặp đôi	Bảng 2×2 trình bày sự phù hợp giữa bệnh và chứng đối với yếu tố nguy cơ Kiểm định χ^2 McNemar OR = tỉ số khả năng phơi nhiễm $= \frac{\text{bãnh cơ chằng khăng}}{\text{bãnh khăng chằng cơ}}$	hồi quy logistic điều kiện
Bất cặp tầng	Phân tích phân tầng: bảng 2×2 cho mỗi tầng Kiểm định χ^2 Mantel-Haenszel $OR = \frac{\Sigma(ad / n)}{\Sigma(bc / n)}$	hồi quy logistic hay phân tích phân tầng
<i>(b) Nhiều chứng cho một bệnh</i>		
Ngẫu nhiên	Như trên	như trên
Một nhóm gồm nhiều bệnh và nhiều chứng của nó	Áp dụng thao tác Mantel-Haenszel đặc biệt	hồi quy logistic có điều kiện
Bất cặp phân tầng	Như trên	như trên
<i>(c) Khoảng tin cậy</i>		
Tất cả các lược đồ	$95\% = OR^{(1 \pm 1,96 / \chi)}$ xấp xỉ dựa vào kiểm định của Miettinen	tính từ hệ số hồi quy

Khoảng tin cậy cho tỉ số số chênh

Khoảng tin cậy cho tỉ số số chênh có thể được tính giống như đối với nguy cơ tương đối, dùng phương pháp dựa vào kiểm định của Miettinen.

$$95\% = OR^{(1 \pm 1,96 / \chi)}$$

Có thể tính được khoảng tin cậy phần trăm khác bằng cách thay 1,96 với điểm phần trăm thích hợp của phân phối bình thường. Công thức này được áp dụng khi kiểm định χ^2 được dùng đối với bảng 2×2 , kiểm định Mantel-Haenszel hay McNemar. Trong cả hai trường hợp tính căn bậc hai của χ^2 .

Thí dụ, xét số liệu bất cặp được trình bày trong bảng 24.5 về sự liên hệ giữa sử dụng thuốc ngừa thai và thuyên tắc mạch, trong đó:

$$OR = 4,4, \chi^2 = 26,4, P < 0,001$$

Do có

$$\chi = (26,4 = 5,14 \text{ và } 1,96/\chi = 1,96/5,14 = 0,38$$

Suy ra

$$OR^{(1+1,96/\chi)} = 4,4^{(1+1,38)} = 4,4^{1,38} = 7,7$$

Và

$$OR^{(1-1,96/\chi)} = 4,4^{(1-1,38)} = 4,4^{0,62} = 2,5$$

Do đó, khoảng tin cậy 95% của tỉ số số chênh là 2,5 tới 7,7.

THỬ NGHIỆM LÂM SÀNG VÀ NGHIÊN CỨU CAN THIỆP

Giới thiệu

Có 3 loại nghiên cứu thực nghiệm (thử nghiệm lâm sàng, vaccine và can thiệp) được giới thiệu ngắn gọn ở Chương 21. Các đặc điểm quy hoạch chính sẽ được mô tả ở đây. Để đơn giản, việc thảo luận sẽ chú trọng và thử nghiệm lâm sàng, nhưng nó cũng được áp dụng đối với thử nghiệm chế độ dự phòng và các biện pháp phòng bệnh khác. Các đặc điểm liên quan đến thử nghiệm vaccine và các can thiệp khác sẽ được xét ngắn gọn ở cuối chương. Xem Pocock (1983) để có thảo luận toàn diện về thử nghiệm lâm sàng.

Thử nghiệm lâm sàng

Một thử nghiệm lâm sàng là một thử nghiệm được tiến hành để đánh giá hiệu quả của một chế độ điều trị mới. Người tình nguyện được phân vào một trong hai nhóm, **nhóm điều trị (treatment group)** và **nhóm chứng (control group)**. **Nhóm chứng** nhận hoặc là một chế độ điều trị chuẩn của bệnh, như trong đánh giá các thuốc mới, hay không điều trị, như trong đánh giá hiệu quả của vaccine. Có thể mở rộng phương pháp bằng cho những nhóm có liều lượng thuốc khác nhau hay có những chế độ điều trị khác nhau.

Cần thiết kế sao cho sự khác biệt giữa nhóm điều trị và chứng có thể quy về tác dụng thực sự của điều trị. Nhóm ban đầu phải giống nhau càng nhiều càng tốt về tất cả các phương diện trừ phương diện trị liệu. Sẽ không hay nếu chẳng hạn như nhóm chứng ban đầu có nhiều người bị bệnh nặng hơn nhóm điều trị, bởi vì chỉ điều này không đã khiến cho trị liệu dường như có hiệu quả. Nó cách khác, phải cẩn thận tránh sai lệch chọn lựa (selection bias). Điều này có thể được bằng cách dùng ngẫu nhiên hóa (randomization), đôi khi được củng cố bởi bắt cặp (matched) hay quy hoạch bắt chéo (cross-over design).

Hơn nữa, phương pháp xử lý và đánh giá các nhóm phải giống nhau. Cơ hội sai lệch trong khi tiến hành thử nghiệm có thể được giảm thiểu bằng cách dùng quy hoạch mù đơn (*single blind*) hay mù đôi (*double-blind*) và placebo.

Ngẫu nhiên hóa

Phân các người tham dự vào nhóm điều trị và nhóm chứng một cách ngẫu nhiên là một cách để tránh các sai lệch chọn lựa (cố tình hay vô ý), hoặc là bởi người nghiên cứu hoặc người tham dự, có thể dẫn đến sự khác nhau giữa hai nhóm. Có thể ngẫu nhiên hóa bằng cách dùng một bảng số ngẫu nhiên (xem Chương 23). Một chữ số được chọn với số lẻ (1,3,5,7,9) tương ứng với nhóm điều trị và số chẵn (0, 2, 4, 6, 8) tương ứng với nhóm chứng (hay ngược lại). Giả sử rằng 6 số đầu tiên được chọn là 2, 7, 9, 0, 4, 7 thì người tham dự thứ nhất sẽ ở nhóm chứng, thứ nhì và thứ ba sẽ ở nhóm điều trị, thứ tư thứ năm sẽ vào nhóm chứng và thứ sáu sẽ vào nhóm điều trị.

Đôi khi người ta muốn phân chia sao cho chẳng hạn như cứ mỗi 10 người tham gia, số tham dự vào mỗi nhóm sẽ bằng nhau. Điều này được gọi là ngẫu nhiên hóa giới hạn (**restricted randomization**) hay ngẫu nhiên hóa cân bằng (randomization with balance). Trong trường hợp này việc sử dụng bảng số ngẫu nhiên sẽ khác đi. Chẳng hạn đối với mỗi 10 người tham gia, bảng được dùng để quyết định 5 người sẽ được phân vào nhóm điều trị bằng cách chọn 5 số ở giữa 01 và 10, giả sử 09, 02, 04, 01, 08 được chọn. Người thứ nhất, thứ hai, thứ tư, tám và thứ chín trong nhóm sẽ được phân vào nhóm điều trị và những người khác, thứ ba, thứ năm, thứ sáu, thứ bảy và thứ mười sẽ được phân vào nhóm đối chứng. quá trình ngẫu nhiên hóa lại được lặp lại cho mỗi 10 người tham gia.

Có thể biến đổi để cho quá trình ngẫu nhiên hóa có thể đảm bảo tỉ lệ nam và nữ trong mỗi nhóm như nhau bằng cách dùng các số ngẫu nhiên riêng cho nam và cho nữ.

Thứ tự sắp xếp tốt nhất phải được quyết định trước khi bắt đầu thử nghiệm. Nên cho một số các phong bì có đánh số ghi rõ sự sắp xếp đã chuẩn bị, được mở mỗi khi một người mới được chấp nhận vào thử nghiệm.

Bắt cặp

Một cách khác, quy hoạch bắt cặp đôi (matched pair design) có thể được dùng để sắp xếp cho nhóm điều trị và nhóm đối chứng tương tự với nhau về các yếu tố gây nhiễu chính, như tuổi và giới tính. Người tham gia được bắt cặp tùy theo biến số bắt cặp, và một người trong mỗi cặp được sắp xếp ngẫu nhiên vào nhóm điều trị và người khác được sắp vào nhóm chứng. Lưu ý rằng việc bắt cặp phải được tuân thủ trong khi phân tích, chẳng hạn như tiến hành kiểm định χ^2 McNemar (xem Chương 14) chứ không phải kiểm định χ^2 2×2 chuẩn.

Quy hoạch bắt cặp đôi không phù hợp bằng ngẫu nhiên hóa khi có sự gia nhập từ từ các người tham gia vào cuộc thử nghiệm tiến hành trong một thời gian dài bởi vì sẽ có sự trì hoãn đáng kể trước khi có thể bắt cặp cho một người mới tham gia.

Quy hoạch bắt chéo

Trong thử nghiệm lâm sàng một chế độ điều trị chỉ có ích lợi tạm thời, có thể dùng chính người tham gia để làm đối chứng. Thí dụ trong khi so sánh 2 thuốc giảm đau (A và B) để điều trị migraine, mỗi người tham gia có thể thử dùng mỗi chế phẩm trong các dịp khác nhau. Thứ tự dùng cần được ngẫu nhiên hóa để cho phân nửa người tham gia dùng thuốc A trước và phân nửa dùng thuốc B trước. Quy hoạch này được gọi là bắt chéo (Cross-over design). Khuyết điểm chính là có thể có tác **động tồn dư (spill-over effect) từ lần thứ nhất ảnh hưởng đến lần thứ hai.**

Quy hoạch mù đơn và mù đôi

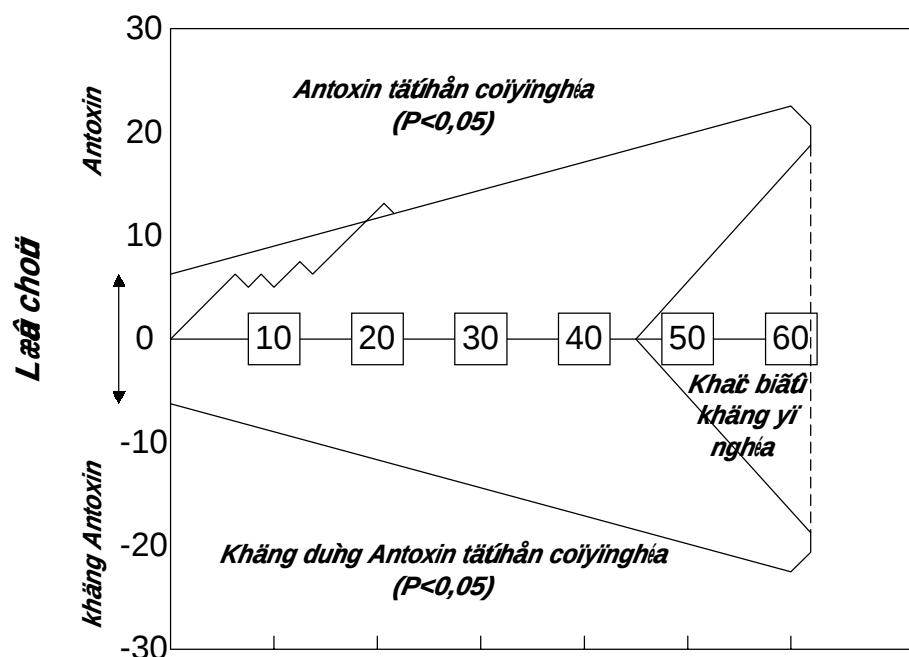
Khi có thể, không nên để người tham gia lẫn người nghiên cứu biết loại trị liệu nào đã sử dụng cho đến cuối cuộc thử nghiệm. Điều này được gọi là quy hoạch mù **đôi** (double blind design) và đảm bảo không bị sai lệch do xử lý và đánh giá nhóm. Điều này đặc biệt cần thiết nếu có sự đánh giá chủ quan. Để đạt được điều này trị liệu cho nhóm chứng cần phải giống về bề ngoài với trị liệu của nhóm điều trị. Điều này có thể là sử dụng chế phẩm placebo không có tác dụng cho nhóm chứng. Chế phẩm điều trị và đối chứng chỉ được nhận biết nhờ mã số, mã số được giữ bởi một người không có tham gia cho đến khi thu thập tất cả các số liệu. Quy hoạch **mù đơn (single blind design)** là khi người nghiên cứu biết nhưng người tham gia không biết về chế độ trị liệu.

Vấn đề đạo đức

Mặc dù về mặt khoa học một thử nghiệm kiểm chứng ngẫu nhiên hóa là cách tốt nhất để xác định ích lợi của một chế độ trị liệu, việc sử dụng chúng đặt ra một vấn đề đạo đức quan trọng. Có đáng phải hoãn trị liệu có thể có ích lợi cho một nhóm đối chứng không? Điều này đặc biệt đúng khi một phương thức trị liệu mới được tin là không gây tác dụng phụ gì và hiện chưa có chế độ điều trị có hiệu quả hay chế độ điều trị chuẩn. Một thí dụ là việc gợi ý bổ sung vitamin trong lúc mang thai có thể giảm số mới mắc của khiếm khuyết ống thần kinh. Cũng vậy, việc nên ngẫu nhiên hóa hay để cho người tham gia có thể chọn nhóm cho họ? có cần tờ cam đoan hay không? Điều này được thảo luận kỹ trong hoàn cảnh của cắt bỏ khối u hay đoạn nhũ để điều trị ung thư vú của Cancer Research Campaign Working Party in Breast Conservation (1983). Để thảo luận đầy đủ hơn vấn đề này và những vấn đề liên quan hãy xem Bradford Hill (1977) trong đó có bản sao Khẳng định của Hội đồng nghiên cứu y khoa về trách nhiệm trong nghiên cứu đối tượng con người (Medical Research Council's Statement on Responsibility in Investigations on Human Subjects) và tuyên ngôn Helsinki (Helsinki Declaration) của Hội đồng y khoa thế giới (World Medical Assembly) về Các khuyến cáo hướng dẫn bác sĩ y khoa trong nghiên cứu y sinh có liên quan đến đối tượng con người (Recommendations Guiding Medical Doctor in Biomedical Research Involving Human Subjects).

Thử nghiệm lâm sàng tuần tự

Thử nghiệm lâm sàng tuần tự (sequential clinical trial) là một thử nghiệm trong đó kết quả được đánh giá liên tục để xem đã có đủ bằng chứng để ngừng thử nghiệm và tuyên bố hoặc là đã xác định được sự khác biệt hay không có sự khác biệt nào.



Hình 25.1 Kết quả từ một thử nghiệm tuần tự để đánh giá giá trị của liều cao kháng độc tố trong điều trị uốn ván lâm sàng. Với sự cho phép của Brown et al. (1960) Lancet ii, 227-230

Thử nghiệm tuần tự thường được dùng khi kết cuộc của đo lường là một tỉ lệ, thí dụ như tỉ lệ sống còn trong mỗi nhóm điều trị. Phương pháp luận được minh họa bằng một thí dụ đặc hiệu (xem chi tiết ở Armitage, 1975). Hình 25.1 trình bày kết quả của một thử nghiệm để đánh giá hiệu quả của liều lượng cao kháng độc tố trong điều trị uốn ván lâm sàng bằng cách so sánh tỉ lệ tử vong giữa những người được và không được cho kháng độc tố. Người tham gia được chia làm 2 nhóm, một người được nhận kháng độc tố và một người không. Nếu người nhận kháng độc tố còn sống trong khi người kia chết sẽ có thiên vị thuận lợi cho kháng độc tố. Ngược lại nếu người nhận kháng độc tố chết còn người kia sống sẽ có thiên vị thuận lợi cho không sử dụng kháng độc tố. Nếu cả hai thành viên đều sống hay đều chết, đây là một cặp phù hợp và không góp thêm thông tin gì cho việc so sánh dùng kháng độc tố với không dùng kháng độc tố.

Kết quả được ghi nhận được trên một đồ thị trình bày số dư của thiên vị và số các cặp không phù hợp. Đồ thị bắt đầu ở điểm zero và vẽ lên một đơn vị và qua phải cho mỗi thiên vị thuận lợi cho kháng độc tố và xuống một đơn vị và qua phải cho mỗi thiên vị thuận lợi cho không dùng kháng độc tố. Ranh giới trên được vẽ để trình bày ý nghĩa ở mức 5% và xác suất đường này hay đường kia bị cắt là 95% nếu tỉ suất của số dư thiên vị là 75% hay hơn. Ở giữa đường hình cái nêm thể hiện không có ý nghĩa. Ranh giới trên bị cắt ở thiên vị thứ 18 và thử nghiệm kết thúc. 15 cặp cho kết quả sống còn thuận lợi cho kháng độc tố và 3 cặp thuận lợi cho không dùng kháng độc tố, do đó xác định sự giảm tử vong có ý nghĩa trong nhóm được cho kháng độc tố ($P < 0,05$).

Ưu điểm chính của thử nghiệm lâm sàng tuần tự là nói chung cần ít người tham gia để có được kết luận hơn một thử nghiệm có kích thước cố định, trong đó phân tích được tiến hành một lần lúc hoàn tất. Dù vậy điều này bù trừ bởi mức ý nghĩa bị giảm bởi vì người ta đã điều chỉnh cho sự gia tăng khả năng có kết quả ý nghĩa do các kiểm định được lặp lại. Do đó, trong thí dụ trên, mức ý nghĩa thu được là 5%. Nếu không có điều chỉnh này, kiểm định McNemar cho giá trị χ^2 là 6,7 khi so sánh 15 cặp thuận lợi với kháng độc tố và chỉ 3 cặp thuận lợi cho việc không dùng kháng độc tố. Nó có mức ý nghĩa 1%. Điều cần xét về mặt thực tiễn là thử nghiệm tuần tự không thích hợp cho các tình huống mà có một khoảng thời gian dài giữa trị liệu và đánh giá cuối cùng của đáp ứng.

Thử nghiệm vaccine

Vấn đề quy hoạch thử nghiệm vaccine tương tự với thử nghiệm lâm sàng thuốc (và các biện pháp điều trị) và nên sử dụng quy hoạch đối chứng ngẫu nhiên hóa. Dù vậy, mục tiêu của thử nghiệm vaccine hoàn toàn khác. Chúng vaccine là một biện pháp phòng bệnh và vaccine và placebo được dùng cho người khỏe mạnh chứ không phải cho người bệnh. Mục tiêu là đánh giá hiệu lực (efficacy) của chủng vaccin trong việc phòng ngừa sự xuất hiện bệnh. Điều này được đo lường bằng tỉ lệ các trường hợp bệnh trong nhóm placebo mà lẽ ra không xảy ra nếu họ được chủng vaccine.

Thí dụ hãy xét bảng 25.1 trong đó trình bày kết quả từ một thử nghiệm một vaccine cúm mới. Tổng số 80 trường hợp cúm xảy ra trong nhóm placebo. Nếu nhóm này được chủng vaccin người ta kì vọng chỉ 8,3% (tỉ suất trong nhóm chủng vaccine) trong số đó bị cúm bằng $220 \times 0,083 = 18,3$. Sẽ cứu được $80 - 18,3 = 61,7$ trường hợp, tính ra kết quả của vaccine là $61,7/80 = 77\%$

Bảng 25.1 Kết quả của thử nghiệm vaccine cúm, đã được trình bày trong bảng 13.1

Cúm	vaccine	placebo	tổng số
Có	20(8,3%)	80(36,4%)	100(21,7%)
Không	220	140	360
Tổng số	240	220	460

Có thể đơn giản hóa tính toán thành công thức sau, cho thấy rõ mối liên kết giữa hiệu lực vaccine và nguy cơ tương đối liên hệ do không sử dụng vaccine.

$$\text{Hiệu lực vaccine (VE)} = 1 - \frac{\text{tỉ suất mắc mào trong nhóm tiêm chủng}}{\text{tỉ suất mắc mào trong nhóm kháng tiêm chủng}}$$

Công thức này cũng được dùng để tính độ tin cậy của thử nghiệm vaccine. Trước hết tính khoảng tin cậy của nguy cơ tương đối như trong Chương 24. Sau đó thay vào công thức trên để có giới hạn tương ứng của hiệu lực vaccine.

$$95\% \text{ c.i. (VE)} = 1 - \frac{1}{RR^{1 \pm 1,96/\chi}}$$

Trong thí dụ này

$$RR = \frac{80/220}{20/240} = 4,36; \quad VE = 1 - \frac{1}{RR} = 1 - \frac{1}{4,36} = 0,77 \text{ (giáng nhẽ trần)}$$

$$\chi^2 = 51,37 \text{ và } \chi = (51,37 = 7,17)$$

Giới hạn tin cậy dưới

$$RR^{1-1,96/7,17} = 4,36^{1-0,27} = 4,36^{0,73} = 2,93$$

$$VE \text{ tương ứng} = 1 - 1/2,93 = 0,66$$

Giới hạn tin cậy trên

$$RR^{1+1,96/7,17} = 4,36^{1+0,27} = 4,36^{1,27} = 6,49$$

$$VE \text{ tương ứng} = 1 - 1/6,49 = 0,85$$

Do đó khoảng tin cậy 95% của hiệu lực vaccine của vaccine cúm này là 66% tới 85%.

Nghiên cứu can thiệp

Thử nghiệm can thiệp là một loại nghiên cứu thực nghiệm khác. Quy hoạch cơ bản của nó tương tự như thử nghiệm lâm sàng chỉ trừ một ngoại lệ. Việc áp dụng của can thiệp thường được áp dụng vào một cộng đồng chứ không phải cho các cá nhân. Thí dụ là đánh giá tác động bổ sung flourur vào nước uống lên bệnh sâu răng ở một số quận này mà không có ở quận khác, và việc đánh giá tác động lên số mới mắc tiêu chảy của một nguồn nước cải thiện. Các vấn đề mà một thử nghiệm dựa trên cộng đồng đặt ra còn là một lãnh vực bị bỏ quên và ngoài phạm vi của cuốn sách này (xem Smith, 1987).

TÍNH CỖ MẪU CẦN THIẾT

Giới thiệu

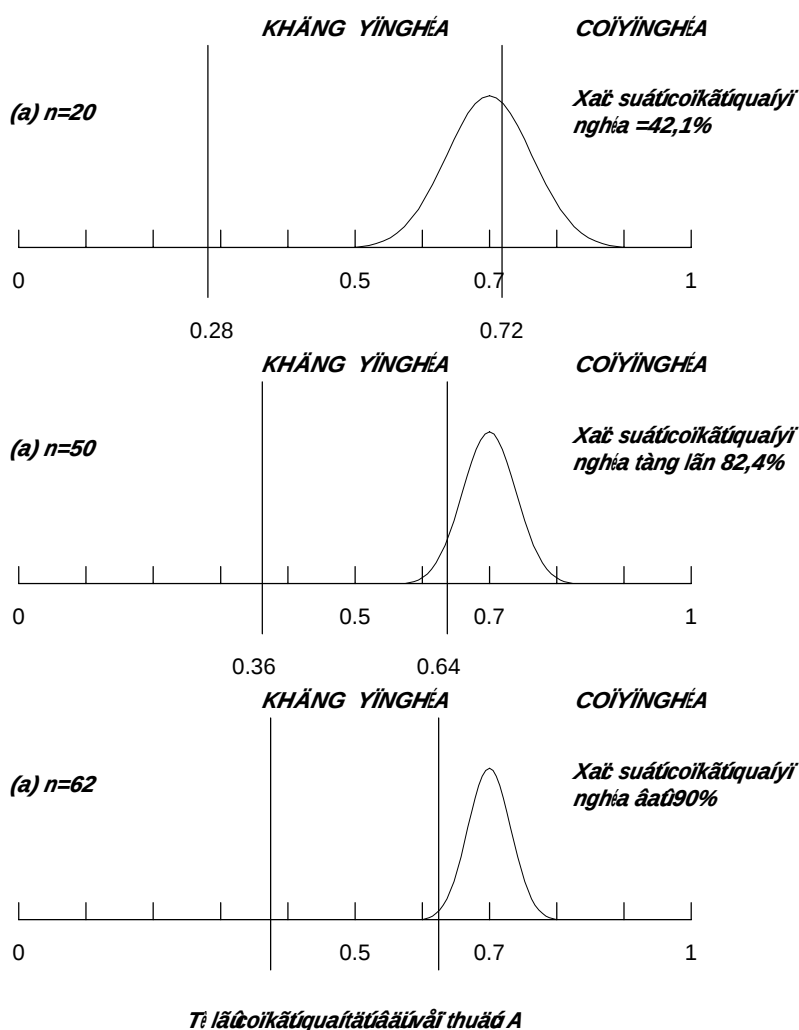
Phần quan trọng trong lập kế hoạch điều tra là quyết định có bao nhiêu người cần được nghiên cứu để trả lời mục tiêu nghiên cứu. Rất nhiều khi người ta quyết định dựa trên lãnh vực hậu cần đơn thuần. Đó là thói quen xấu và nhiều người xem đó là không đạo đức. Mặt khác nghiên cứu nhiều người hơn cần thiết là sự lãng phí thời gian, tiền bạc và tài nguyên (thường là giới hạn). Trong thử nghiệm lâm sàng điều này không chỉ là nhiều người hơn cần thiết phải sử dụng placebo mà còn có nghĩa là làm trì trệ việc đưa vào sử dụng trị liệu có ích. Mặt khác tham gia một nghiên cứu quá nhỏ không trả lời được mục tiêu cũng đáng chê trách không kém.

Dù vậy, tính toán cỡ mẫu cần thiết không phải là một thao tác đơn giản bởi vì trước hết cần phải định lượng hóa mục tiêu. Thí dụ, không thể chỉ khẳng định đơn thuần rằng mục tiêu là xác định có phải trẻ nuôi bằng bình có nguy cơ tử vong cao hơn trẻ nuôi bằng sữa mẹ hay không. Nó cũng cần khẳng định độ lớn của nguy cơ gia tăng muốn xác định bởi vì để phát hiện sự gia tăng gấp 4 lần cần cỡ mẫu nhỏ hơn khi muốn phát hiện sự gia tăng gấp 2 lần. Hơn nữa, ngay cả khi cỡ mẫu gia tăng, bởi vì sự biến thiên lấy mẫu (xem Chương 3), luôn luôn có khả năng chỉ quan sát được nguy cơ nhỏ hơn rất nhiều và chúng ta không thể đảm bảo rằng nghiên cứu cho kết quả ý nghĩa. Chúng ta có thể làm gì để định rõ xác suất mà chúng ta muốn đạt được một ý nghĩa thống kê nhất định. Xác suất này gọi là năng lực (power) của nghiên cứu. Do đó chúng ta có thể xác định rằng nghiên cứu sẽ xứng đáng nếu có xác suất 90% khẳng định được sự khác biệt trong nguy cơ giữa nhóm trẻ bú bình và bú mẹ nếu nguy cơ tương đối thực sự là 2.

Nguyên lý của việc xác định cỡ mẫu

Trong chương 12 chúng ta đã thảo luận việc phân tích kết quả của thử nghiệm lâm sàng để so sánh 2 loại thuốc giảm đau, trong đó 9 trong số 12 người bị migraine nói rằng họ giảm nhiều với thuốc A hơn thuốc B. Dù vậy, tỉ lệ vượt trội 0,75 thuận lợi cho thuốc A không có ý nghĩa ($P=0,15$), cho thấy rằng kết quả phù hợp giới giả thuyết trung tính là không có sự khác biệt về hiệu quả giữa hai loại thuốc. Có thể xảy ra một trong hai điều hoặc là thực sự không có sự khác biệt giữa các loại thuốc hoặc là thuốc A tốt hơn nhưng cỡ mẫu 12 không đủ lớn để cho thấy điều đó.

Xét quy hoạch một nghiên cứu như vậy, và chúng ta cần bao nhiêu bệnh nhân để có thể xác định tính chất ưu việt của thuốc A. Đầu tiên chúng ta phải đặc hiệu hóa ưu việt hơn nghĩa là gì. Chúng ta bắt đầu bằng cách khẳng định điều này là tỉ lệ thiên lệch 70% hay hơn về thuốc A và chúng ta quyết định chúng ta muốn năng lực 90% đạt được mức ý nghĩa 5%.



Hình 26.1 Xác suất đạt được ý nghĩa (ở mức 5%) với các cỡ mẫu khác nhau (n) khi kiểm định tỉ lệ thiên vị thuốc A hơn thuốc B khác giá trị giả thuyết trung tính là 0,5, nếu giá trị thực là 0,7

Nguyên lí của việc xác định cỡ mẫu được giới thiệu tốt nhất bằng cách xem xét các cỡ mẫu khác nhau và đánh giá chúng có thỏa mãn yêu cầu của chúng ta hay không. Chúng ta sẽ bắt đầu với cỡ mẫu 20, được mô tả trong hình 26.1 (a). Nhớ lại rằng, kết quả có mức ý nghĩa 5% nếu nó cách trung bình 1,96 s.e. Hay hơn, và giá trị của giả thuyết trung tính về tỉ lệ thiên lệch với thuốc A là 0,5 và sai số chuẩn (s.e.) bằng $\sqrt{0,5 \times 0,5/n}$. Điều này có nghĩa là cỡ mẫu bằng 20?

$$s.e. = \sqrt{0,5 \times 0,5 / 20} = 0,1118$$

$$0,5 + 1,96 s.e. = 0,5 + 1,96 \times 0,1118 = 0,72$$

$$\text{và } 0,5 - 1,96 s.e. = 0,5 - 1,96 \times 0,1118 = 0,28$$

Do đó phạm vi giá trị có kết quả ý nghĩa là trên hoặc bằng 0,72 và dưới hoặc bằng 0,28.

Nếu tỉ lệ thực sự là 0,7, khả năng quan sát được tỉ lệ bằng hoặc lớn hơn 0,72 và cho kết quả có ý nghĩa là bao nhiêu? Khả năng này được minh họa bởi hình gạch chéo trên hình. Đường cong thể hiện phân phối lấy mẫu, đó là phân phối bình thường có sai số chuẩn là $(0,7 \times 0,3/20) = 0,1025$. Giá trị z tương ứng với 0,72 là:

$$\frac{0,72 - 0,7}{0,1025} = 0,20$$

Tỉ lệ của phân phối bình thường lớn hơn 0,20 được tìm thấy trong bảng A1 bằng 0,421 hay 42,1%. Do đó, nói tóm lại điều này có nghĩa là với cỡ mẫu 20 ta chỉ có 41,2% cơ hội xác định được thuốc A tốt hơn nếu tỉ lệ thiên vị thực sự là 0,7.

Sau đó chúng ta hãy xem chuyện gì xảy ra nếu ta tăng cỡ mẫu lên 50, như trình bày trong hình 26.1(b). Phạm vi của giá trị mở rộng đến từ 0,64 lên trên và từ 0,36 xuống dưới. Phân phối lấy mẫu đã hẹp lại và vùng trùng lấp với phạm vi ý nghĩa lớn hơn. Do đó xác suất có kết quả ý nghĩa đã tăng lên bằng 82,5% nhưng chưa đạt đến yêu cầu của chúng ta là 90%.

Do đó chúng ta nhất định sẽ cần phải hơn 50 bệnh nhân để có năng lực 90%. Nhưng chúng ta cần bao nhiêu? Chúng ta cần tăng cỡ mẫu tới n , ở đó vùng trùng lấp giữa phân phối lấy mẫu và phạm vi ý nghĩa đạt được 90% như trình bày trong hình 26.1(c). Chúng ta sẽ mô tả cách làm thế nào để tính trực tiếp cỡ mẫu cần thiết. Đạt được kết quả có ý nghĩa nếu chúng ta quan sát được giá trị lớn hơn

$$0,5 + 1,96 \text{ s.e.} = 0,5 + 1,96 \times \sqrt{(0,5 \times 0,5/n)}$$

(hay dưới $0,5 - 1,96 \text{ s.e.}$). Chúng ta muốn chọn n đủ lớn để 90% phân phối lấy mẫu ở trên điểm này. Giá trị z của phân phối lấy mẫu tương ứng với 90% là - 1,28 (xem bảng A2), có nghĩa là một giá trị quan sát

$$0,7 - 1,28 \text{ s.e.} = 0,7 - 1,28 \times \sqrt{(0,7 \times 0,3/n)}$$

Do đó n được chọn đủ lớn để

$$0,7 - 1,28 \times \sqrt{(0,7 \times 0,3/n)} > 0,5 + 1,96 \times \sqrt{(0,5 \times 0,5/n)}$$

Chuyển về ta được

$$0,7 - 0,5 > \frac{1,96 \times \sqrt{0,5 \times 0,5} + 1,28 \times \sqrt{0,7 \times 0,3}}{\sqrt{n}}$$

Bình phương hai vế và chuyển về ta được

$$\begin{aligned} n &> \frac{[1,96 \times \sqrt{0,5 \times 0,5} + 1,28 \times \sqrt{0,7 \times 0,3}]^2}{0,2^2} \\ &= \frac{1,5666^2}{0,2^2} = 61,4 \end{aligned}$$

Do đó chúng ta cần khoảng 60 bệnh nhân để thỏa mãn yêu cầu của chúng ta là có năng lực 90% xác định sự khác biệt có ý nghĩa giữa thuốc A và thuốc B ở mức 5%, nếu sự thiên lệch thực sự cho thuốc A là 0,7.

Công thức tính cỡ mẫu

Thảo luận ở trên liên quan đến việc xác định cỡ mẫu cho phép kiểm định một tỉ lệ đơn, đó là tỉ lệ người tham dự thiên lệch thuốc A so với thuốc B. Trên thực tiễn không nhất thiết lúc nào cũng phải lí luận chi tiết như vậy. Cỡ mẫu có thể tính trực tiếp từ công thức tổng quát, dùng cho kiểm định ý nghĩa một tỉ lệ đơn là

$$n > \frac{\left\{ u\sqrt{\pi(1-\pi)} + v\sqrt{\pi_0(1-\pi_0)} \right\}^2}{(\pi - \pi_0)^2}$$

Trong đó

- n = cỡ mẫu cần thiết
- p = tỉ lệ quan tâm
- p₀ = tỉ lệ của giả thuyết trung tính
- u = điểm phần trăm một đuôi của phân phối bình thường tương ứng với 100% - năng lực, thí dụ nếu năng lực là 90%, (100% - năng lực) = 10% và u = 1,28
- v = điểm phần trăm hai đuôi của phân phối bình thường tương ứng mức ý nghĩa cần thiết, thí dụ nếu mức ý nghĩa là 5%, v = 1,96

Áp dụng công thức này cho thí dụ trên ta có:

$$p = 0,7 \quad p_0 = 0,5 \quad u = 1,28 \quad \text{và} \quad v = 1,96$$

Suy ra

$$n > \frac{[1,28\sqrt{0,7 \times 0,3} + 1,96\sqrt{0,5 \times 0,5}]^2}{(0,7 - 0,5)^2}$$

$$= \frac{1,5666^2}{0,2^2} = 61,4$$

Cũng đúng bằng kết quả thu được ở trên.

Cũng có thể áp dụng các nguyên lí tương tự cho các trường hợp khác. Ở đây chúng ta không lí luận chi tiết nhưng sẽ trình bày các công thức dùng cho các tình huống phổ biến nhất trong bảng 26.1. Danh sách gồm 2 phần. Bảng 26.1(a) gồm các trường hợp mà mục đích nghiên cứu là xác định sự khác biệt có ý nghĩa. Bảng 26.1(b) gồm các tình huống ước lượng một con số quan tâm với độ chính xác nhất định. Lưu ý rằng trong trường hợp có hai trung bình, hai tỉ lệ hay hai tỉ suất công thức cho cỡ mẫu cần thiết cho một trong hai nhóm. Do đó tổng số cỡ mẫu cần thiết cho nghiên cứu gấp đôi số này. Bảng 26.2 trình bày các hệ số điều chỉnh cho việc *quy hoạch nghiên cứu có các nhóm không cùng cỡ* (xem thí dụ 26.3). Lưu ý rằng trong công thức áp dụng cho tỉ suất cỡ mẫu có cùng đơn vị với tỉ suất (xem thí dụ 26.2)

Việc sử dụng bảng sẽ được minh họa qua một số ví dụ. Cần nhận thức rằng việc tính toán cỡ mẫu chỉ dựa trên sự suy đoán tốt nhất của chúng ta về tình hình. Con số không phải là con số ảo thuật. Nó chỉ cho một ý niệm về độ lớn của số người cần được nghiên cứu. Nói cách khác, nó hữu dụng khi cần phân biệt giữa 50 và 100 chứ không phải để phân biệt giữa 51 và 52. Vì lí do này, các hiệu chỉnh liên tục không được đưa vào công thức liên quan đến tỉ lệ, bởi vì nó được coi như sự phức tạp không cần thiết. Hơn nữa, các suy đoán của chúng ta khó có thể chính xác, bởi vì nếu chúng ta đã biết câu trả lời chúng ta đã chẳng cần làm nghiên cứu. Do đó cần tiến hành tính cỡ mẫu cho một số hoàn cảnh chứ không phải chỉ trong một hoàn cảnh. Như vậy sẽ có một hình ảnh rõ ràng về phạm vi có thể của cuộc nghiên cứu và rất hữu ích trong việc cân đối giữa cái mong muốn và tính khả thi về hậu cần.

Cuối cùng, cỡ mẫu phải tăng để bù trừ cho những người không đáp ứng và bỏ cuộc. Cũng cần điều chỉnh khi sử dụng lược đồ lấy mẫu cụm chứ không phải lấy mẫu ngẫu nhiên, và cần thiết kiểm soát yếu tố gây nhiễu trong phân tích. Điều này không chỉ ngoài phạm vi cuốn sách mà còn là một lĩnh vực chưa mấy phát triển.

Thí dụ 26.1

Một nghiên cứu được tiến hành trong một vùng nông thôn ở Đông Phi để xem việc bổ sung thức ăn trong lúc mang thai có làm tăng trọng lượng trẻ sinh ra hay không. Phụ nữ khám tiền

sản được phân ngẫu nhiên thành có nhận thức ăn bổ sung và không nhận thức ăn bổ sung. Công thức 4 trong bản 26.1 sẽ giúp chúng ta quyết định cần huy động bao nhiêu phụ nữ trong mỗi nhóm. Chúng ta cần cung cấp những thông tin sau:

(a) độ lớn của sự khác nhau về cân nặng mà chúng ta muốn phát hiện. chúng ta quyết định rằng sự gia tăng 0,25 kg là tác động đáng kể mà chúng ta không muốn bỏ qua. Do đó chúng ta áp dụng công thức này với $\mu_1 - \mu_2 = 0,25$ kg.

Bảng 26.1 Công thức xác định cỡ mẫu (a) cho các nghiên cứu xác định sự khác biệt có ý nghĩa (b) cho các nghiên cứu nhằm ước lượng một chỉ số quan tâm với độ chính xác nhất định

		Thông tin cần thiết	Công thức tính cỡ mẫu tối thiểu
(a) Kết quả ý nghĩa			
1. Trung bình đơn	$\mu - \mu_0$	Sự khác biệt giữa trung bình, μ với giá trị giả thuyết trung tính	$\frac{\mu_1 + \mu_2}{e^2}$
		Độ lệch chuẩn	
	σ	Xem ở dưới	
2. Tỉ suất đơn*	u, v		$\frac{(u + v)^2 \mu}{(\mu - \mu_0)^2}$
	μ	Tỉ suất	
	μ_0	Giá trị giả thuyết trung tính	
3. Tỉ lệ đơn	u, v		$\frac{\left\{ u\sqrt{\pi(1-\pi)} + v\sqrt{\pi_0(1-\pi_0)} \right\}^2}{(\pi - \pi_0)^2}$
	π	Tỉ lệ	
	π_0	Giá trị giả thuyết trung tính	
4. So sánh hai trung bình (cỡ mẫu cho mỗi nhóm)	u, v		$\frac{(u + v)^2 (\sigma_1^2 + \sigma_2^2)}{(\mu_1 - \mu_2)^2}$
	$\mu_1 - \mu_2$	Hiệu số hai trung bình	
	σ_1, σ_2	Độ lệch chuẩn	
5. So sánh hai tỉ suất* (cỡ mẫu cho mỗi nhóm)	u, v		$\frac{(u + v)^2 (\mu_1 + \mu_2)}{(\mu_1 - \mu_2)^2}$
	μ_1, μ_2	Tỉ suất	
		Xem ở dưới	
6. So sánh hai tỉ lệ (cỡ mẫu cho mỗi nhóm)	u, v		$\frac{\left\{ u\sqrt{\pi_1(1-\pi_1)} + \pi_2(1-\pi_2) + v\sqrt{2\pi(1-\pi)} \right\}^2}{(\pi_2 - \pi_1)^2}$
	π_1, π_2	Tỉ lệ	
		Xem ở dưới	
7. Nghiên cứu bệnh chứng (cỡ mẫu cho mỗi nhóm)	π_1	Tỉ lệ nhóm chứng tiếp xúc	$\pi_2 = \frac{\pi_1 OR}{1 + \pi_1 (OR - 1)}$
	OR	Tỉ số số chênh	
	π_2	Tỉ lệ nhóm bệnh tiếp xúc	
Tất cả các trường hợp	u, v	Xem ở dưới	
	u	Điểm phần trăm một đuôi của phân phối bình thường tương ứng với 100%-năng lực, td. nếu năng lực = 90%, $u = 1,28$	
	v	Điểm phần trăm của phân phối bình thường tương ứng với mức ý nghĩa, td. nếu mức ý nghĩa = 5%, $v = 1,96$	
(b) Độ chính xác			
8. Trung bình đơn	σ	Độ lệch chuẩn	$\frac{\sigma^2}{e^2}$
	e	Độ lớn sai số chuẩn cần thiết	
9. Tỉ suất đơn*	μ	Tỉ suất	$\frac{\mu}{e^2}$
	e	Độ lớn sai số chuẩn cần thiết	
10. Tỉ lệ đơn	π	Tỉ lệ	$\frac{\pi(1-\pi)}{e^2}$
	e	Độ lớn sai số chuẩn cần thiết	
11. Hiệu số giữa hai trung bình (cỡ mẫu cho mỗi nhóm)	σ_1, σ_2	Độ lệch chuẩn	$\frac{\sigma_1^2 + \sigma_2^2}{e^2}$
	e	Độ lớn sai số chuẩn cần thiết	
12. Hiệu số giữa hai tỉ suất* (cỡ mẫu cho mỗi nhóm)	μ_1, μ_2	Tỉ suất	$\frac{\mu_1 + \mu_2}{e^2}$
	e	Độ lớn sai số chuẩn cần thiết	
11. Hiệu số giữa hai tỉ lệ (cỡ mẫu cho mỗi nhóm)	π_1, π_2	Tỉ lệ	$\frac{\pi_1(1-\pi_1) + \pi_2(1-\pi_2)}{e^2}$
	e	Độ lớn sai số chuẩn cần thiết	

* Trong các trường hợp này, cỡ mẫu có cùng đơn vị với mẫu số của tỉ suất. Thí dụ nếu tỉ suất tính theo người năm thì công thức tính ra người năm quan sát (xem thí dụ 26.2)

(b) Độ lệch chuẩn của phân phối trọng lượng lúc sinh trong mỗi nhóm. Chúng ta quyết định giả sử rằng độ lệch chuẩn của trọng lượng lúc sinh sẽ như nhau trong hai nhóm. Số liệu cũ gợi ý rằng độ lệch chuẩn là 0,4 kg. Nó khác chúng ta giả sử rằng $s_1 = 0,4$ kg và $s_2 = 0,4$ kg.

(c) Thống nhất năng lực cần thiết là 95%. Do đó chúng ta cần $u = 1,64$

(d) Mức ý nghĩa cần thiết. Chúng ta quyết định rằng nếu có thể chúng ta sẽ đạt được mức ý nghĩa 1%. Do đó $v = 2,58$. Áp dụng công thức 4 với các giá trị này ta đạt được

$$\begin{aligned} n &> \frac{(1,64 + 2,58)^2 \times (0,4^2 + 0,4^2)}{0,25^2} \\ &= \frac{17,8084 \times 0,32}{0,0625} = 91,2 \end{aligned}$$

Do để thỏa mãn yêu cầu của chúng ta chúng ta cần huy động khoảng 90 người phụ nữ trong mỗi nhóm.

Thí dụ 26.2

Trước khi tham gia vào một can thiệp cung cấp nước, thanh khiết và vệ sinh ở nam Bangladesh, đầu tiên chúng ta muốn biết số lần tiêu chảy hằng năm của trẻ dưới 5 tuổi. Chúng ta giả sử tỉ suất mới mắc là 3 nhưng muốn ước lượng chúng với sai số (0,2). Điều này có nghĩa là, nếu chúng ta quan sát được 2,6 lần tiêu chảy/trẻ/năm chúng ta có thể kết luận rằng tỉ suất thực sự có lẽ ở giữa 2,4 và 2,8 lần/trẻ/năm. Viết lại điều này bằng từ ngữ thống kê, chúng ta muốn rằng khoảng tin cậy 95% không rộng hơn (0,2). Bởi vì độ rộng của khoảng tin cậy này khoảng (2 s.e.). Điều này có nghĩa là chúng ta muốn nghiên cứu đủ trẻ để có sai số chuẩn nhỏ hơn 0,1 lần/trẻ/năm. Áp dụng công thức 9 và bảng 26.1 tính ra được

$$n > \frac{3}{0,1^2} = 300$$

Lưu ý rằng công thức này áp dụng cho tỉ suất (số 3, 6, 11, 14) cho số mẫu cần thiết có đơn vị giống như tỉ suất. Chúng ta tính tỉ suất theo số lần cho mỗi trẻ cho mỗi năm. Do đó chúng ta cần nghiên cứu 300 trẻ năm để có độ chính xác mong muốn. Điều này có thể đạt được bằng cách quan sát 300 trẻ trong một năm hay quan sát 1200 trẻ trong 3 tháng. Dù vậy không nên bỏ qua các khả năng khác như tác động của thời tiết khi quyết định khoảng thời gian của nghiên cứu.

Thí dụ 26.3

Một nghiên cứu bệnh chứng để nghiên cứu xem trẻ bú bình có nguy cơ chết vì nhiễm trùng hô hấp cấp tính có cao so với trẻ bú mẹ hay không. Mẹ ở một nhóm bệnh (trẻ chết với giấy chứng tử có nguyên nhân chết là hô hấp sẽ được phỏng vấn về tình trạng bú sữa mẹ của trẻ trước khi chết. Kết quả được so sánh với số liệu thu được về tình trạng bú sữa mẹ từ mẹ của những đứa trẻ chứng bình thường. Người ta kì vọng khoảng 40% chứng ($p_1 = 0,4$) sẽ bú bình và chúng ta muốn phát hiện sự khác biệt nếu bú bình liên quan đến sự gia tăng nguy cơ tử vong gấp 2 ($OR = 2$). Cần nghiên cứu bao nhiêu bệnh và chứng để có năng lực 90% ($u = 1,28$) đạt được mức ý nghĩa 5% ($v = 1,96$)? việc tính toán gồm một số bước chi tiết trong công thức 7 của bảng 26.1.

Bảng 26.2 Hệ số điều chỉnh để sử dụng trong quy hoạch nghiên cứu để so sánh các nhóm có kích thước không bằng nhau, như việc chọn nhiều chứng cho một bệnh trong nghiên cứu bệnh chứng. Hệ số (f) áp dụng cho nhóm nhỏ hơn và bằng $(c+1)/2c$, trong đó kích thước của nhóm lớn hơn sẽ bằng c lần kích thước của nhóm nhỏ. Cỡ mẫu của nhóm nhỏ sẽ là fn trong đó n là số cần thiết trong trường hợp 2 nhóm bằng nhau và cỡ mẫu của nhóm lớn sẽ là cnf (xem thí dụ 26.3)

Tỉ số của nhóm lớn so với nhóm nhỏ (c)	Điều chỉnh cỡ mẫu của nhóm nhỏ (f)
1	1
2	3/4
3	2/3
4	5/8
5	3/5
6	7/12
7	4/7
8	9/16
9	5/9
10	11/12

(a) Tính p_2 , tỉ lệ của trẻ bú bình

$$\pi_2 = \frac{\pi_1 OR}{1 + \pi_1(OR - 1)} = \frac{0,4 \times 2}{1 + 0,4 \times (2 - 1)} = \frac{0,8}{1,4} = 0,57$$

(b) Tính ($\bar{p}$, trung bình của p_1 và p_2)

$$\bar{\pi} = \frac{0,4 + 0,57}{2} = 0,485$$

(c) Tính cỡ mẫu tối thiểu

$$n > \frac{[1,28\sqrt{0,4 \times 0,6 + 0,57 \times 0,43} + 1,96\sqrt{2 \times 0,485 \times 0,515}]^2}{(0,57 - 0,4)^2}$$

$$= \frac{[1,28 \times \sqrt{0,4851} + 1,96 \times \sqrt{0,4996}]^2}{0,17^2} = \frac{2,2769^2}{0,17^2} = 179,4$$

Do đó chúng ta cần huy động khoảng 1 bệnh và 180 chứng để có tổng số cỡ mẫu là 360. Chuyện gì sẽ xảy ra nếu chúng ta quyết định huy động số chứng lớn gấp 3 lần số bệnh. Bảng 26.2 trình bày các hệ số điều chỉnh cho số các bệnh tùy theo số các chứng/bệnh khác nhau. Thí dụ $c = 3$, hệ số điều chỉnh là $2/3$. Điều này có nghĩa là chúng ta cần $180 \times 2/3 = 120$ bệnh và số chứng nhiều gấp 3 lần = 360. Mặc dù số chứng cần thiết đã giảm đáng kể, cỡ mẫu chung đã tăng từ 360 lên 540.

SỬ DỤNG MÁY TÍNH

Giới thiệu

Trong thập kỉ qua đã diễn ra cuộc cách mạng trong sự tiếp cận với công nghệ máy tính. Mười lăm năm trước, rất ít người có thể có được thậm chí một máy tính cầm tay và đa số các máy tính cầm tay đều có thành tích tầm thường, chỉ làm các phép tính số học cơ bản. Máy tính là các máy lớn cần một môi trường được điều hòa không khí và một số người điều hành chúng. Chúng được sử dụng giới hạn trong các nhà phân tích viên máy tính và các lập trình viên. Sự phát triển của máy vi tính đã làm thay đổi tất cả. Với một chi phí tương đối thấp, có thể mua cho mình một hệ thống máy tính đặt vừa trên một chiếc bàn và thỏa mãn hầu hết các nhu cầu tính toán và học được (ngay cả khi trước đó không có kinh nghiệm về máy tính) các kĩ năng về tính toán trong thời gian tương đối ngắn.

Chỉ trừ các nghiên cứu rất nhỏ, hầu hết các nghiên cứu đều có ích khi sử dụng máy tính. Nó tính toán nhanh hơn và chính xác hơn khi tính bằng tay. Có thể xử lí số liệu lớn và phân tích phức tạp hơn. Nó cũng cho phép dễ dàng lập lại các phân tích có thay đổi, chẳng hạn như phân tích nam và nữ riêng và chung. Dù vậy cần phải lưu ý một số điều. Việc sử dụng máy tính đòi hỏi thêm thì giờ để chuẩn bị số liệu và chương trình. Cần phải kiểm tra tất cả các thủ tục trên một số nhỏ các số liệu và kiểm chứng kết quả bằng tay. Nếu không có thể xảy ra các sai sót nghiêm trọng. Cuối cùng, có thể đi rất ra và làm rất nhiều các phân tích mà không có suy nghĩ thực sự về chúng.

Chúng ta sẽ mô tả ngắn gọn các phương tiện căn bản cho bất kì các hệ thống máy tính nào và xác định các hạn chế phổ biến. Bởi vì công nghệ máy tính phát triển rất nhanh, chúng tôi không khuyến cáo nhất định về loại thiết bị cần thiết, và danh sách cũng không toàn diện. Để nắm thêm chi tiết cần một cuốn sách chuyên ngành cập nhật hay một cuốn sách hướng dẫn cho người mua máy tính.

Phần cứng máy tính

Thiết bị máy tính được gọi là phần cứng (hardware). Một máy vi tính điển hình gồm có một bàn phím, một màn hình (monitor), một bộ vi xử lí có bộ nhớ (bộ nhớ truy cập ngẫu nhiên - Random access memory - RAM), và các ổ đĩa. Nó thường được kèm theo máy in.

Một máy tính lớn (mainframe) có bộ xử lí mạnh hơn, có hệ thống chia sẻ thời gian giải quyết với nhiều người dùng cùng một lúc. Nó có hệ điều hành tinh vi hơn và có khả năng lưu trữ lớn hơn. Số liệu có thể ở trên đĩa cứng để dễ dàng và nhanh chóng truy suất trong khi phân tích hay trên các băng từ để lưu trữ số lượng lớn trong thời gian lâu dài. Dung lượng của băng từ phụ thuộc vào chiều dài, mật độ và dạng thức. Thí dụ, một băng từ dài 2400 feet được viết với mật độ 1600 b.p.i. (bits per inch) có thể giữ khoảng 30 triệu kí tự thông tin. Nếu nó chứa một bản văn lục dài 120 biến số, có tổng số 300 chữ số, băng từ sẽ có thể giữ khoảng 100.000 bản văn lục.

Ổ đĩa

Có hai loại ổ đĩa, đĩa mềm và đĩa cứng. Ổ đĩa mềm có thể tiếp nhận đĩa mềm có kích thước khác nhau (thí dụ 3,5 inch, 5,25 inch) phụ thuộc vào máy cụ thể. Dung lượng của đĩa (và của bộ nhớ trong) được tính bằng byte, trong đó một byte là số không gian cần thiết để chứa đựng một kí tự. Khả năng của đĩa phụ thuộc vào dạng thức của hệ thống sử dụng để sắp xếp số liệu trên đĩa, thông tin được lưu trữ mật độ đơn hay mật độ đôi, và đĩa có thể đọc một mặt hay hai mặt. Chữ k kí hiệu 100 và mega có nghĩa là một triệu. Do đó một đĩa mềm có dung lượng 1,44 megabyte có thể chứa được 1.400.000 kí tự.

Đĩa cứng thường có dạng cố định, nằm ở bên trong máy tính. Nó có dung lượng lớn chẳng hạn như 120 hay 210 megabyte, truy suất đọc và ghi rất nhanh.

Tổ chức dữ liệu

Số liệu được lưu trữ dưới dạng tập tin (file), và được tập hợp trong thư mục (directory) của đĩa. Một tập tin có những mẫu tin (records). Có thể có một hay nhiều mẫu tin cho một cá nhân chẳng hạn tương ứng với số trang khác nhau của bản vấn lục. Tập tin có thể có cấu trúc tuần tự hay truy cập ngẫu nhiên. Trong tập tin **tuần tự (sequential file) các mẫu tin được viết và đọc theo thứ tự. Thí dụ, để truy cập** thông tin có liên quan đến người số 10, đầu tiên cần đọc thông tin từ người số 1 đến 9. Tập tin truy cập ngẫu nhiên (random access file) có thể đi thẳng đến vị trí của mẫu tin số 10 để đọc (hay ghi) thông tin cho người số 10.

Các mẫu tin được viết dưới dạng thức cố định hay tự do. Trong dạng thức cố **định** (fixed format), số liệu được sắp xếp theo chuỗi các số (hay kí tự), và một biên luôn luôn ở cùng một vị trí cột trong mẫu tin. Số được căn phải (right justified). Thí dụ, nếu hai vị trí được phân cho tuổi, thì 9 tuổi sẽ được lưu trữ bằng 09 hay 9 chứ không phải là 9. Dấu thập phân không được ghi rõ; vị trí của nó được rút ra từ số cột. Trong dạng thức tự do (free format), biến số được sắp xếp theo mẫu tin, cách nhau bởi dấu phẩy hay khoảng trắng, không cố ý xếp hàng theo mẫu tin. Phải ghi các số thập phân. Một thí dụ của hai loại dạng thức trên được trình bày trong bảng 27.1

Cuối cùng số liệu có thể trữ trực tiếp dưới dạng kí tự (dạng ascii) hay được chuyển sang dạng máy gồm một loạt các số 0 và 1 (dạng nhị phân - binary format). Các dạng thức được dùng phụ thuộc vào tập tin đã được tạo ra như thế nào. Nhiều phần mềm sử dụng dạng nhị phân cho bất kì mọi thao tác và xử lí số liệu trong phần mềm nhưng chúng tạo ra tập tin dưới dạng thức ascii nếu nó sử dụng với các phần mềm khác hay chuyển tới máy khác.

Bảng 27.1 So sánh dạng thức cố định và tự do để lưu trữ số liệu

ID	Tuổi	số con	tăng trọng (kg)	trọng lượng trẻ sơ sinh (g)
1	21	1	11,5	3520
2	18	0	9,5	2950
3	25	1	10,6	3380
4	31	3	14,1	3860
5	23	0	9,9	3200
Dạng thức cố định			Dạng thức tự do	
12111153520			1,21,1,11.5,3520	
21800952950			2,18,0,9.5,2950	
32511063380			3,25,1,10.6,3380	
43131413860			4,31,3,14.1,3860	
5230 993200			5,23,0,9.9,3200	

Sao chép lưu

Cần phải có ít nhất 3 bản sao của bất kì tập tin số liệu nào. Cả đĩa và băng từ có thể bị hư hay bị phá hủy. Tốt nhất nên giữ ít nhất một bản sao ở một nơi khác để an toàn không bị các biến cố như cháy, lụt lội và đánh cắp. Bản sao chép lưu (back-up copy) phải được thực hiện thường xuyên. Điều này không bao giờ được lưu ý đầy đủ.

Hơn nữa, khi được nhập vào máy tính, số liệu thường được lưu ngay vào bộ nhớ máy tính và chỉ được chuyển vào đĩa khi bộ nhớ đầy. Tốt nhất là phải đảm bảo rằng số liệu được chuyển vào đĩa thường xuyên, thí dụ như mỗi nửa giờ, để giảm thiểu số liệu bị mất và lãng phí thời gian khi máy tính bị dừng vì một số lí do. Mất điện hay treo máy do lỗi, thí dụ như cố gắng

ghi nhiều thông tin vào đĩa không còn trống, hay vô tình nhấn một số phím vào đó và có thể mất số liệu trong bộ nhớ và có thể gây phá hủy đĩa.

Phần mềm máy tính

Không có máy tính nào có thể làm việc được mà không có phần mềm (soft ware). Ở mức thấp nhất là hệ điều hành (operating system), gồm các lệnh cần thiết để điều hành máy tính như ghi chép (COPY) tập tin, ghi ra màn hình (TYPE) một tập tin hay thực hiện (RUN) một chương trình, liệt kê thư mục (DIR) của một đĩa. Kế đến là **chương trình (program) gồm một bộ lệnh để làm một nhiệm vụ nào đó, thí dụ** như so sánh trọng lượng của trẻ với trọng lượng trung vị của tuổi đó ở trong tiêu chuẩn phát triển và tính trọng lượng quan sát được theo phần trăm của chuẩn. Một chương trình được viết trong một ngôn ngữ (language), thí dụ như Basic, Fortran, C, hay Pascal. Chương trình được chuyển thành lệnh máy nhờ bộ biên dịch (compiler). Cuối cùng nhiều chương trình áp dụng chung được gọi là gói phần **mềm (package)**.

Các gói phần mềm mà chúng ta quan tâm có 3 loại, căn cứ dữ liệu (database), thống kê và xử lý văn bản (word-processing). Phần mềm căn cứ dữ liệu nên dùng để nhập số liệu và tạo các cấu trúc dữ liệu tổng quát, như sắp xếp và chọn một nhóm các mẫu tin. Hầu hết sẽ tạo các bản vấn lục trên màn hình và hiện lên tên các biến được nhập. Con trỏ (cursor) dùng để chỉ vị trí hiện tại và chuông sẽ kêu khi hoàn thành bản vấn lục. Lưu ý rằng một số gói phần mềm có hạn chế về tổng số biến có thể nhập cho một cá nhân nhưng có thể khắc phục bằng cách kết hợp các biến, chẳng hạn như tất cả các biến nhân trắc thành một biến dài.

Nhiều chương trình căn cứ dữ liệu cũng có thể lập bảng và tính các thống kê cơ bản. Dù vậy nói chung nên dùng gói phần mềm thống kê (statistical package) chuyên dụng để thao tác và phân tích số liệu. Phần mềm xử lý văn bản (word-processing package) hữu ích không chỉ để báo cáo kết quả mà còn để xây dựng chương trình và quy hoạch bản vấn lục. Liên kết với gói phần mềm căn cứ dữ liệu, có thể dùng để tạo ra bản vấn lục đã có một phần thông tin như tên, địa chỉ, tuổi và tránh chép bằng tay từ cuộc điều tra này sang cuộc điều tra khác.

Dù vậy, đôi khi các đòi hỏi của cuộc điều tra không thể thỏa mãn bằng bộ phần mềm. Do đó cần phải viết một số chương trình nhỏ.

CHÈ MUÛC

95% confidence limits..... 25

A

a priori.....	48
Âa giaïc tảo suát.....	10
âaĩnh dáu.....	9
ân vê báuc hai.....	15
ân vê báuc mẩu.....	151
ăăă the.....	5
ăăăng biăău săă.....	64
ăăăng nhăăt	
phăăn phăăi.....	11
Ăăă âăăc hiăău.....	144
ăăă dăăc.....	54
ăăă khai dă căăc âăăi.....	93
ăăă lăăch bệnh thăăng chăăă.....	21
ăăă lăăch chăăă.....	14
tênh toăăi.....	16
ăăă lăăch chăăă hệnh hoăăc.....	125
Ăăă nhăăy căăm.....	144
ăăă tăăi do	
cúă phăăng sai.....	14
ăăăi xăăng.....	11
additive rule.....	65

 $\mathcal{A}$ Æāic læāüng..... 135

A

ảnh lễ giải hân trung tâm.....	26
ảnh tên.....	4
agesex adjusted mortality rate.....	99
aiêm 5 phảo trăm.....	24
aiêm châu.....	54
aiêm phảo trăm.....	23
aiêu tra cảo ngang	
nhiều lần.....	137
analysis of variance.....	41
ANCOVA.....	<i>See</i> phân tẻch iểu phẩng sai
ANOVA.....	41
ao iáo.....	180
ÄÖ iáo cẩng.....	180
ÄÖ iáo mẩm.....	180
Äo lẩng tẩ vong	
cẩ sẩ iáo.....	97
trong nghiẩ cẩ.....	96
aoain hẩ.....	138
arithmetic mean.....	13
association.....	135
attributable risk.....	155
attributable risk percent.....	155
äuäi dẩẩ.....	11
äuäi trẩ.....	11

 B

backup copy	181
Bảng 2 × 2 (so sánh hai tế lâu)	77
bảng 2 × c	
cảng thực ruít gỏin	82
bảng đầu trui	77

baíng đæù trui 2×2	77
baíng gáúp 4	77
Baíng lăín (nhiăôu haíng, nhiăôu căüt)	80
baíng săú ngăúu nhiăín	148
baíng săúng	105
so săính	107
baíng săúng áoain hăú	105
baíng săúng hiăúín hainh	105
balanced design	45
balanced with replication	45
bàòng phán phăúí bệnh thăểíng	20
băôt căúp táóng	159
bar diagram	7
băú biăín dệch	182
băúnh (case)	138
băúnh hiăúm	
giăi âệnh	157
băúnh phăô biăúín	
giăi âệnh	158
Bayes	65
Bayesian approach	65
Berkson's fallacy	143
between groups	42
bias	139, 142
biăúín âăôi logistic	93
biăúín săú	4
biăúín săú gáy nhiăúu	89
biăúín săú nhệ phán	68
biăúín thiăín	14
biăúín thiăín láúy mắúu	4
bimodal	11
binary format	180
binary variable	68
binomial distribution	68

C

cải màu	4, 170
nguyên lệ xác định	170
Cảng thực tên cải màu	173
cảng nằm thanh	9
case control study	138
casecontrol study	158
câu hỏi àoing	140
Câu hỏi mắi	140
censored	105
census	148
central limit theorem	26
chằng trệh	181
chặing (control)	138
chỉ square test	77
childyears at risk	98
Choãn màu hầi thắing	149
Choãn màu ngầi hầi ắi	148
clinical trials	139
close question	140
cluster	150, 152
cluster analysis	64
cluster sampling	152
<i>Coefficient of variation</i>	16
cohort	138
compiler	182
components of variation	50
Con trối	182

conditional logistic regression	161	false negative rate	144
conditioned logistic regression	94	false positive rate	144
confidence interval	19	file	180
confounding bias	143	final digit	119
consistency checks	141	finite population correction	19
contingency table	77	first stage units	151
continous monitoring	137	fivebar gates	9
continuity correction	74	fixed	138
control group	164	fixed effects	50
correlation	51	fourfold table	77
correllation coefficient	52	free format	180
correspondance analysis	64	frequencies	7
covariate	64	relative	7
Cox's propotional model	108	frequency	
cross product ratio	158	distribution	8
crossover design	164	frequentist definition	65
Crossover design	165	Friedman twóway ANOVA	130
crossectional study	137		
cumulative	106	<i>G</i>	
current life table	105	Gaussian distribution	20
cursor	182	Generalized Linear Interactive Modelling	94
cũm	152	geometric mean	123
		geometric standard deviation	125
<i>D</i>		geomtric mean	14
dán sǎu		giái thuyǎut trung tênh	30
Phán phǎui tǎon suǎut	11	giǎi trẽ bẽ khuyǎut	50
Dán sǎu	3	giǎi hǎun tin cáuy	25
dán sǎu câu ǎnh	138	giǎim sǎit liǎn tuúc	137
dán sǎu ǎũng	138	gõii phǎon mǎom	182
dán sǎu muúc tiǎu	4	gõii phǎon mǎom thǎung kǎ	182
database	182		
daũng ascii	180	<i>H</i>	
daũng thǎic câu ǎnh	180	hai ǎuǎi	24
degrees of freedom	14, 27	hai yǎu vẽ	11
dependent	66	hǎim phán biǎut	64
dependent variable	54	hǎoi cǎu	137, 153
deviance	94	Hǎoi quy bæǎic luii	62
deviation	14	Hǎoi quy bæǎic tǎii	62
diǎũn tẽch dǎũn ǎũng cong	22	hǎoi quy bǎui	42, 58
digit	148	vai phán tẽch phǎng sai	63
direct method	99	Hǎoi quy bǎui vǎn 2 biǎũn sǎu	59
directory	180	Hǎoi quy bǎui vǎn cǎic biǎũn giǎi thẽch phi tuyǎũn tênh	63
discordant pairs	160	Hǎoi quy bǎui vǎn cǎic biǎũn giǎi thẽch rǎi raũc62	
discriminant function	64	Hǎoi quy bǎui vǎn nhiǎũu biǎũn	61
discriminat analysis	64	hǎoi quy logistic	64, 93
double blind design	166	hǎoi quy logistic cõ ǎiũũ kiǎũn	94, 161
doubleblind controlled trial	139	Hǎoi quy tǎo hǎũp tǎui ǎu	62
dummy variable	62	Hǎoi quy tuyǎũn tênh	54
dynamic	138	Hǎoi quy vǎo trung bẽnh	145
		hardware	179
<i>E</i>		hǎu ǎiũũ hǎinh	181
efficacy	168	hǎu sǎu hǎoi quy riǎng phǎon	60
equal probability scheme	137	hǎu sǎu tǎang quan	52
equal probability selection method epsem	150	vai báing phán tẽch phǎng sai	59
estimation	135	hǎu sǎu tǎang quan bǎui	60
evaluation	136	hazard rate	108
exact test	79, 84	Helsinki Declaration	166
experimental	136	hiǎũu chẽnh tênh liǎn tuúc	74
explanatory variable	54	hiǎũu chẽnh tênh liǎn tuúc cuía Yates	78
exponential	108	hiǎũu lǎũc	168
		hiǎũu sǎu hai trung bẽnh	
<i>F</i>		phán phǎui lǎũy mǎũu	37
factor	41	hierarchical structure	151
factor analysis	64	histogram	9
factorial	70		

I

ì thành phần của sự biến thiên.....	50
incidence.....	97
incidence rate.....	98
incidence risk.....	98
independent variable.....	54
indirect method.....	99
information bias.....	143
interaction.....	45
intercept.....	54
intervention trials.....	139

J

J ngữ cảnh	
phân phối.....	11

K

kế hoạch sàng.....	107
kê tẩu.....	180
Kendall's rank correlation.....	130
khảo sát tin cậy.....	19, 28
khung mẫu.....	148
Kiểm định bệnh tật.....	34
kiểm định chênh lệch.....	79
Fisher.....	84
Kiểm định chênh lệch	
cho bảng 2 x 2.....	84
kiểm định chi bệnh tật.....	77
Kiểm định chi bệnh tật.....	77
Kiểm định chi bệnh tật.....	91
kiểm định log rank.....	107
Kiểm định mất mát.....	33
Kiểm định phụ thuộc	
phân phối chi bệnh tật.....	118
Kiểm định sắp xếp hạng coi dấu Wilcoxon.....	131
Kiểm định t cặp đôi.....	30
kiểm định t hai đầu.....	33
Kiểm định t mất mát.....	33
Kiểm định tổng sắp xếp hạng Wilcoxon.....	132
kiểm định ý nghĩa	
quan hệ với khảo sát tin cậy.....	32
Kiểm định ý nghĩa.....	30
đúng xác định bệnh tật.....	73
kiểm định.....	141
Kiểm tra phân vị.....	141
Kiểm tra sáu lần.....	141
Kiểm tra tên phụ thuộc.....	141
Kolmogorov-Smirnov test.....	130
Kruskal Wallis oneway ANOVA.....	130

L

Lưu ý gia.....	136
language.....	181
lưu.....	11
lưu mẫu.....	148
lưu mẫu.....	150, 152
Lưu mẫu.....	152
lưu mẫu hai bước.....	151
lưu mẫu hoàn lại.....	151
lưu mẫu nhiều bước.....	150, 152
Lưu mẫu nhiều bước.....	151
lưu mẫu phân tầng.....	150
Lưu mẫu phân tầng.....	150
least squares fit.....	54
liên tục.....	4

life expectancy.....	107
life table.....	105
linear.....	51
linear regression.....	51
log linear model.....	94
logistic regression.....	93
logit.....	127
lognormal distribution.....	122
longitudinal study.....	137

M

mã hệ số.....	108
mã hệ số nguy cơ tử vong của Cox.....	108
mã hệ số tuyến tính log.....	94
Mã thức sàng lọc.....	108
mức ý nghĩa.....	30, 31
main effect.....	45
mainframe.....	179
máy tính	
trong phân tích thống kê.....	179
máy tính cá nhân.....	6
trong hồi quy tuyến tính.....	57
máy tính lần.....	179
Mann Whitney U test.....	130
mã số toàn bộ.....	97
mã số toàn bộ thời gian.....	97
matched.....	164
matched data.....	94
matched design.....	160
matched pair design.....	165
matched paired design.....	86
mất mát.....	24
mất mát về.....	11
mẫu.....	3, 148
mẫu cái lần.....	25
Mẫu mới.....	26
mẫu tin.....	180
maximum likelihood.....	93
mean square.....	42
median.....	13
method of linear contrasts.....	48
missing value.....	50
mode.....	13
monitor.....	179
MS regression.....	59
MS residual.....	59
multiple invasion.....	120
multiplicative rule.....	65
multiple comparison.....	79
multiple correlation coefficient.....	60
multiple regression.....	42
multiple regression equation.....	59
multiple response questions.....	141
multi-stage sampling.....	152
multistage.....	150
multistage sampling.....	151

N

năng lực.....	35, 170
natural logarithms.....	123
ngăn ngừa.....	181
ngẫu biến Berkson.....	143
ngẫu nhiên hóa cân bằng.....	165
ngẫu nhiên hóa giải phẫu.....	165
nguyên nhân.....	

lấp kầu hoaých	135	phân phôi Poisson	
muối tiêu	135	Hệ thống	111
tiểu hân	135	Phân phôi Poisson	110
nguyên cứu áoan hân	153	ánh nhê	110
Nghiên cứu bân chêng	138, 158	vai tề suất	113
Nghiên cứu can thiáp	169	phân phôi t	27
nguyên cứu cật dúc	137	vân (n-1) âu tâu do	27
Nghiên cứu cật ngang	137	phân phôi tón suất	
Nghiên cứu quan sát	136	Hệ thống	11
nguyên cứu thêi nghiâm		phân sâu lúy màu	19, 150
bật chêu	165	phân sâu quy trạch dân sâu	156
Nghiên cứu thêi nghiâm	139	Phân tân âô	52
nguy cả		phân tềc âa biân	64
bân hiâm - bân phâo biân	98	Phân tềc âa biân	64
Nguy cả mât môt	98	phân tềc âông phâng sai	63
Nguy cả qui trạch	154	phân tềc cuôm	64
nguy cả quy trạch	155	Phân tềc nhiâu bâng 2×2	88
nguy cả quy trạch dân sâu	156	phân tềc phâng sai	41
nguy cả tâng âu	99, 154	trong hân quy tuyền tênh	58
nhân cûp khâng phui hân	160	yâu tâu	41
nhê phán	4, 180	Phân tềc phâng sai hai chiâu	45
nhôm âiâu trê	164	Phân tềc phâng sai mât chiâu	42
nhôm chêng	164	Phân tềc phân biân	64
normal distribution	20	Phân tềc sâu liâu	5
null hypothesis	30	phân tềc sâng côi	105
numerical taxonomy	64	phân tềc tâng âng	64
<i>O</i>		Phân tềc tề suất	104
observational	136	Phân tềc thân phân chên	64
odds	156	Phân tềc thâng kâ hân tềc	136
odds ratio	156	Phân tềc yâu tâu	64
onê sided	24	phân cêng	179
onesided	33	phân môm	181
oneway analysis of variance	42	Phân môm can cêi dâi liâu	182
open question	140	Phân môm mât tênh	181
operating system	181	Phân môm xêi lê vân bân	182
Optimal combination regression	62	phân trui	4
overall sample value	150	phân vi	14
<i>P</i>		phêi biân âô	
package	182	lêu chôn	127
paired	86	Phêi biân âô	122
paired t test	30	logarithm	122
partial correlations	60	phêi biân âô lúy cên	127
partial regression coefficients	60	phêi biân âô lúy lúy thêi	127
period prevalence	97	phêi biân âô nghêc âiô	127
persons at risk	97	phi tham sâu	28
personyears at risk	98	phâng phâi	129
personyears of observation	96	pie chart	7
phâng phâi chôn lêu xêi suất bân nhau	150	placebo	166
phâng phâi chuân hoai giân tiáp	99	point prevalence	97
phâng phâi chuân hoai trêc tiáp	99	population attributable fraction	156
Phâng sai	14	population attributable risk	156
phân loâi hêc sâ	64	population proportional attributable risk	156
phân phôi		posterior probability	65
kiôm ánh tênh phui hân	116	power	35, 170
phân phôi bênh thêi		power transformation	127
kiôm ánh tênh phui hân	116	powerful	41
phân phôi bênh thêi chuân	21	prevalence	97
Phân phôi bênh thêi chuân	21	pricipal component analysis	64
phân phôi F	43	prior probability	65
phân phôi nhê thêi	68	probability proportional to size p.p.s	151
xáp xêi phân phôi bênh thêi	73	probit	116
Phân phôi nhê thêi	68	probit plot	116
		program	181
		prophylactic trials	139
		proportion	68
		proportional attributable risk	155

prospective	137, 153
<i>Q</i>	
quan sát	4, 136
Quy hoạch cân ngẫu cỡ mẫu	45
Quy hoạch bán vắn lức	139
quy hoạch bất cập ngẫu	165
quy hoạch cân ngẫu kháng cỡ mẫu	45
Quy hoạch cân ngẫu cỡ mẫu	45
Quy hoạch cân ngẫu kháng mẫu	46
quy hoạch kháng cân ngẫu	45
Quy hoạch kháng cân ngẫu	48
quy hoạch mũi ngẫu	166
Quy hoạch mũi ngẫu	166
Quy tắc cũng	65
xác suất	66
Quy tắc nhân	65
xác suất	65
<i>R</i>	
rối loạn	4
RAM	179
random access file	180
random effects	50
randomization	164
randomization with balance	165
randomized controlled doubleblind design	142
range	14
range checks	141
rank	131
rate	95
reciprocal transformation	127
records	180
Registrar General	136
regression coefficient	54
relative risk	99, 154
Relative risk	154
repeated cross sectional surveys	137
residual sum of squares	42, 58
restricted randomization	165
restrospective	137
retrospective	153
reverse Jshaped	11
risk	96
robustness	27
root transformation	127
<i>S</i>	
sả ngẫu xác suất cân bằng	137
sâu liên quan	135
sâu sai lệch khỏi ảnh hưởng phụ mẫu	92
sai lầm	
trong kiểm định giả thuyết	35
sai lệch	89, 142
Sai lệch gây nhiễu	143
Sai lệch thăng tin	143
sai số	
hệ thống	142
ngẫu nhiên	142
nguồn gốc	142
Sai số chuẩn mẫu	142
sai số chuẩn	17
của trung bình mẫu	17

Sai số lấy mẫu	17
sample	148
sampling fraction	19, 150
sampling frame	148
Sao chép lại	181
sâu chỉnh	156
sâu lũy tích	117
Sâu mũi mũi	97
sâu mũi toàn bộ thải khoáng	97
sâu ngẫu nhiên quan sát	96
scatter diagram	52
screening test	144
secondstage units	151
selection bias	142, 164
sensitivity	144
sequential clinical trial	166
sequential file	180
sign test	130
significance level	30
significant test	30
simple random sampling	148
single blind design	166
skewed	11
So sánh 2 tẻ mẫu	86
so sánh 2 trung bình	
mẫu lớn hay biểu ngẫu lệch chuẩn	38
so sánh hai trung bình	
Cải mẫu nhỏ, ngẫu lệch chuẩn kháng bằng nhau	40
mẫu lớn hay biểu ngẫu lệch chuẩn	37
So sánh hai trung bình	37
So sánh nhiều trung bình... See phân tích phương sai	
soft ware	181
Spearman	133
specificity	144
spillover effect	165
standard deviation	14
standard error of the sample mean	17
standard normal deviate	21
standard normal distribution	21
standard population	99
standardization	99
standardized mortality ratio	99
statistical package	182
Step down regression	62
Stepup regression	62
strata	150
stratified	150
study design	135
subjective definition	65
sum of square	42
summary test	89
symmetrical	11
systematic	142
<i>T</i>	
t distribution	27
table of random numbers	148
tảng phân tuyến tính	48
Tảng quan	52
tảng quan riêng phần	60
Tảng quan sấp hướng Spearman	133
tảng quan tích moment Pearson	133
tảng tích	45

taïc ãũũg cãũ ãẽnh	50	khoãng tin cãũy	25
taïc ãũũg chẽnh	45	tẽnh toãin	16
taïc ãũũg ngãũu nhiãn	50	trung bẽnh bẽnh phãẽng	42
taïc ãũũg tãũn dẽ	165	trung bẽnh chuãũn hoĩa	100
tails		trung bẽnh nhãn	14, 123
upper and lower	11	trung vẽ	13
tallying	9	tuyãn ngãn Helsinki	166
Tãũ chãẽc ããũ	9	two stage sampling	151
Tãũ chãẽc dẽi liãũu	180	two sided	24
tãũn suãũt		twosided	33
sãũ liãũu ãẽnh læũũg	8		
Tãũn suãũt		<i>U</i>	
sãũ liãũu ãẽnh tẽnh	7	unbalanced design	45
tãũn suãũt tãẽng ããũi	7	uniform	11
tãũng	150	unimodal	11
tãũng ãiãũu tra	148		
tãũng bẽnh phãẽng hãũi quy	58	<i>V</i>	
tãũng bẽnh phãẽng phãũn dẽ	42, 58	vaccine trials	139
tãũng cãic bẽnh phãẽng	42	validity	79
tãũng cãic bẽnh phãẽng ãẽũũc giãĩi thẽch bãĩi hãũi quy	58	variance	14
tãũp tin	180	varianceratio	43
Tãũp tin truy cãũp ngãũu nhiãn	180	verification	141
tãũp tin tuãũn tãũ	180	vital statistics	136
Tẽ lãũ	68		
tẽ lãũ ããn		<i>W</i>	
kiãũm ãẽnh yĩ nghẽa	71	Wilcoxon	131
tẽ sãũ sãũ chãẽnh	156	Wilcoxon rank sum test	132
tẽ sãũ tãĩ vong chuãũn hoĩa	99	Wilcoxon sign rank test	130, 131
tẽ suãũt ãm tẽnh giãĩ	144	with replacement	151
Tẽ suãũt chuãũn hoĩa	99	with replication	45
tẽ suãũt dẽũng tẽnh giãĩ	144	within groups	42
tẽ suãũt mãĩi mãũc	98	without replication	45
phãn tẽch	113	word-processing package	182
tẽ suãũt nguy hãũi	108	wordprocessing	182
Tẽ suãũt sinh	95		
Tẽ suãũt tãĩ vong	95	<i>X</i>	
tẽ suãũt tãĩ vong ãẽũũc ãiãũu chẽnh theo tuãũi giãĩi tẽnh	99	xaĩc suãũt	65
Tẽ suãũt tãĩ vong ãẽũũc ãiãũu chẽnh theo tuãũi giãĩi tẽnh	103	ãẽnh nghẽa chui quan	65
tẽnh vãĩng	27	ãẽnh nghẽa theo tãũn suãũt	65
thãe muóc	180	quy tãũc	65
Thãĩ nghiãũm can thiãũp	139	xaĩc suãũt ãiãũu kiãũn	66
thãĩ nghiãũm cõi kiãũm soãĩt mui ããĩ	139	xaĩc suãũt hãũu nghiãũm	65
Thãĩ nghiãũm lãũm saĩng	139, 164	xaĩc suãũt tẽ lãũ vãĩi kẽch thãẽĩc	151
thãĩ nghiãũm saĩng loũc	144	xaĩc suãũt tiãũn nghiãũm	65
Thãĩ nghiãũm vaccine	139, 168		
thãũc nghiãũm	136	<i>Y</i>	
thãĩi gian sãũng cõin	105	Yates' continuity correction	78
Thãũng kã	3	yãũu vẽ	13
Thay ããũi ããn vẽ	17		
tiãũn cãĩũ	137, 153	<i>Z</i>	
trẽnh bãĩy kãũt quãĩ	5	z 21	
treatment group	164		
trung bẽnh	13	<i>χ</i>	
		χ^2 goodness of fit	130

Danh sách các bảng số thống kê

Bảng A1 Diện tích ở đuôi của phân phối chuẩn bình thường	2
Bảng A2 Điểm phần trăm của phân phối bình thường chuẩn	3
Bảng A3 Điểm phần trăm của phân phối t	4
Bảng A4 Điểm phần trăm của phân phối F	5
Bảng A5 Điểm phần trăm của phân phối χ^2	8
Bảng A6 Probits	9
Bảng A7 Giá trị tới hạn của kiểm định sắp hạng có dấu cặp đôi Wilcoxon	11
Bảng A8 Giá trị tới hạn của kiểm định tổng sắp hạng Wilcoxon.....	12
Bảng A9 Giá trị tới hạn của hệ số tương quan sắp hạng của Spearman	16
Bảng A10 Số ngẫu nhiên.....	17

Bảng A1 Diện tích ở đuôi của phân phối chuẩn bình thường

Cải biên từ Bảng 3 của White et al. có xin phép tác giả và nhà xuất bản

z	số lẻ thứ nhì của z									
	0.00	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.0	0.5000	0.4960	0.4920	0.4880	0.4840	0.4801	0.4761	0.4721	0.4681	0.4641
0.1	0.4602	0.4562	0.4522	0.4483	0.4443	0.4404	0.4364	0.4325	0.4286	0.4247
0.2	0.4207	0.4168	0.4129	0.4090	0.4052	0.4013	0.3974	0.3936	0.3897	0.3859
0.3	0.3821	0.3783	0.3745	0.3707	0.3669	0.3632	0.3594	0.3557	0.3520	0.3483
0.4	0.3446	0.3409	0.3372	0.3336	0.3300	0.3264	0.3228	0.3192	0.3156	0.3121
0.5	0.3085	0.3050	0.3015	0.2981	0.2946	0.2912	0.2877	0.2843	0.281	0.2776
0.6	0.2743	0.2709	0.2676	0.2643	0.2611	0.2578	0.2546	0.2514	0.2483	0.2451
0.7	0.2420	0.2389	0.2358	0.2327	0.2296	0.2266	0.2236	0.2206	0.2177	0.2148
0.8	0.2119	0.2090	0.2061	0.2033	0.2005	0.1977	0.1949	0.1922	0.1894	0.1867
0.9	0.1841	0.1814	0.1788	0.1762	0.1736	0.1711	0.1685	0.1660	0.1635	0.1611
1.0	0.1587	0.1562	0.1539	0.1515	0.1492	0.1469	0.1446	0.1423	0.1401	0.1379
1.1	0.1357	0.1335	0.1314	0.1292	0.1271	0.1251	0.1230	0.1210	0.1190	0.1170
1.2	0.1151	0.1131	0.1112	0.1093	0.1075	0.1056	0.1038	0.1020	0.1003	0.0985
1.3	0.0968	0.0951	0.0934	0.0918	0.0901	0.0885	0.0869	0.0853	0.0838	0.0823
1.4	0.0808	0.0793	0.0778	0.0764	0.0749	0.0735	0.0721	0.0708	0.0694	0.0681
1.5	0.0668	0.0655	0.0643	0.0630	0.0618	0.0606	0.0594	0.0582	0.0571	0.0559
1.6	0.0548	0.0537	0.0526	0.0516	0.0505	0.0495	0.0485	0.0475	0.0465	0.0455
1.7	0.0446	0.0436	0.0427	0.0418	0.0409	0.0401	0.0392	0.0384	0.0375	0.0367
1.8	0.0359	0.0351	0.0344	0.0336	0.0329	0.0322	0.0314	0.0307	0.0301	0.0294
1.9	0.0287	0.0281	0.0274	0.0268	0.0262	0.0256	0.0250	0.0244	0.0239	0.0233
2.0	0.02275	0.02222	0.02169	0.02118	0.02068	0.02018	0.0197	0.01923	0.01876	0.01831
2.1	0.01786	0.01743	0.017	0.01659	0.01618	0.01578	0.01539	0.01500	0.01463	0.01426
2.2	0.01390	0.01355	0.01321	0.01287	0.01255	0.01222	0.01191	0.01160	0.01130	0.01101
2.3	0.01072	0.01044	0.01017	0.0099	0.00964	0.00939	0.00914	0.00889	0.00866	0.00842
2.4	0.00820	0.00798	0.00776	0.00755	0.00734	0.00714	0.00695	0.00676	0.00657	0.00639
2.5	0.00621	0.00604	0.00587	0.0057	0.00554	0.00539	0.00523	0.00508	0.00494	0.00480
2.6	0.00466	0.00453	0.0044	0.00427	0.00415	0.00402	0.00391	0.00379	0.00368	0.00357
2.7	0.00347	0.00336	0.00326	0.00317	0.00307	0.00298	0.00289	0.0028	0.00272	0.00264
2.8	0.00256	0.00248	0.0024	0.00233	0.00226	0.00219	0.00212	0.00205	0.00199	0.00193
2.9	0.00187	0.00181	0.00175	0.00169	0.00164	0.00159	0.00154	0.00149	0.00144	0.00139
3.0	0.00135	0.00131	0.00126	0.00122	0.00118	0.00114	0.00111	0.00107	0.00104	0.00100
3.1	0.00097	0.00094	0.0009	0.00087	0.00084	0.00082	0.00079	0.00076	0.00074	0.00071
3.2	0.00069	0.00066	0.00064	0.00062	0.00060	0.00058	0.00056	0.00054	0.00052	0.00050
3.3	0.00048	0.00047	0.00045	0.00043	0.00042	0.00040	0.00039	0.00038	0.00036	0.00035
3.4	0.00034	0.00032	0.00031	0.00030	0.00029	0.00028	0.00027	0.00026	0.00025	0.00024
3.5	0.00023	0.00022	0.00022	0.00021	0.00020	0.00019	0.00019	0.00018	0.00017	0.00017
3.6	0.00016	0.00015	0.00015	0.00014	0.00014	0.00013	0.00013	0.00012	0.00012	0.00011
3.7	0.00011	0.00010	0.00010	0.00010	0.00009	0.00009	0.00008	0.00008	0.00008	0.00008
3.8	0.00007	0.00007	0.00007	0.00006	0.00006	0.00006	0.00006	0.00005	0.00005	0.00005
3.9	0.00005	0.00005	0.00004	0.00004	0.00004	0.00004	0.00004	0.00004	0.00003	0.00003

Bảng A2 Điểm phần trăm của phân phối bình thường chuẩn

Giá trị P	Điểm phần trăm	
	Một bên	Hai bên
0.5	0.00	0.67
0.4	0.25	0.84
0.3	0.52	1.04
0.2	0.84	1.28
0.1	1.28	1.64
0.05	1.64	1.96
0.02	2.05	2.33
0.01	2.33	2.58
0.005	2.58	2.81
0.002	2.88	3.09
0.001	3.09	3.29
0.0001	3.72	3.89

Bảng A3 Điểm phần trăm của phân phối t

Cải biên từ Bảng 7 của White et al. có xin phép tác giả và nhà xuất bản

d.f.	P một bên								
	0.25	0.1	0.05	0.025	0.01	0.005	0.0025	0.001	0.0005
	P hai bên								
	0.5	0.2	0.1	0.05	0.02	0.01	0.005	0.002	0.001
1	1.00	3.08	6.31	12.71	31.82	63.66	127.32	318.29	636.58
2	0.82	1.89	2.92	4.30	6.96	9.92	14.09	22.33	31.60
3	0.76	1.64	2.35	3.18	4.54	5.84	7.45	10.21	12.92
4	0.74	1.53	2.13	2.78	3.75	4.60	5.60	7.17	8.61
5	0.73	1.48	2.02	2.57	3.36	4.03	4.77	5.89	6.87
6	0.72	1.44	1.94	2.45	3.14	3.71	4.32	5.21	5.96
7	0.71	1.41	1.89	2.36	3.00	3.50	4.03	4.79	5.41
8	0.71	1.40	1.86	2.31	2.90	3.36	3.83	4.50	5.04
9	0.70	1.38	1.83	2.26	2.82	3.25	3.69	4.30	4.78
10	0.70	1.37	1.81	2.23	2.76	3.17	3.58	4.14	4.59
11	0.70	1.36	1.80	2.20	2.72	3.11	3.50	4.02	4.44
12	0.70	1.36	1.78	2.18	2.68	3.05	3.43	3.93	4.32
13	0.69	1.35	1.77	2.16	2.65	3.01	3.37	3.85	4.22
14	0.69	1.35	1.76	2.14	2.62	2.98	3.33	3.79	4.14
15	0.69	1.34	1.75	2.13	2.60	2.95	3.29	3.73	4.07
16	0.69	1.34	1.75	2.12	2.58	2.92	3.25	3.69	4.01
17	0.69	1.33	1.74	2.11	2.57	2.90	3.22	3.65	3.97
18	0.69	1.33	1.73	2.10	2.55	2.88	3.20	3.61	3.92
19	0.69	1.33	1.73	2.09	2.54	2.86	3.17	3.58	3.88
20	0.69	1.33	1.72	2.09	2.53	2.85	3.15	3.55	3.85
21	0.69	1.32	1.72	2.08	2.52	2.83	3.14	3.53	3.82
22	0.69	1.32	1.72	2.07	2.51	2.82	3.12	3.50	3.79
23	0.69	1.32	1.71	2.07	2.50	2.81	3.10	3.48	3.77
24	0.68	1.32	1.71	2.06	2.49	2.80	3.09	3.47	3.75
25	0.68	1.32	1.71	2.06	2.49	2.79	3.08	3.45	3.73
26	0.68	1.31	1.71	2.06	2.48	2.78	3.07	3.43	3.71
27	0.68	1.31	1.70	2.05	2.47	2.77	3.06	3.42	3.69
28	0.68	1.31	1.70	2.05	2.47	2.76	3.05	3.41	3.67
29	0.68	1.31	1.70	2.05	2.46	2.76	3.04	3.40	3.66
30	0.68	1.31	1.70	2.04	2.46	2.75	3.03	3.39	3.65
40	0.68	1.30	1.68	2.02	2.42	2.70	2.97	3.31	3.55
60	0.68	1.30	1.67	2.00	2.39	2.66	2.91	3.23	3.46
120	0.68	1.29	1.66	1.98	2.36	2.62	2.86	3.16	3.37
vô cực	0.67	1.28	1.64	1.96	2.33	2.58	2.81	3.09	3.29

Bảng A4 Điểm phần trăm của phân phối F

Cải biên từ Bảng 4 của Armitage (1971) và bảng 18 của Pearson & Harley (1966) có xin phép tác giả và nhà xuất bản và Biometrika Trustees

Các bảng trình bày kiểm định một bên để so sánh 2 phương sai, được dùng trong phân tích phương sai. Tính kiểm định hai bên bằng cách nhân đôi giá trị P

d.f.₁ = d.f. của tử số; d.f.₂ = d.f. của mẫu số

d.f. ₂	P	d.f. ₁														
		1	2	3	4	5	6	7	8	9	10	20	40	60	120	vô cực
1	0.05	161	199	216	225	230	234	237	239	241	242	248	251	252	253	254
1	0.025	648	799	864	900	922	937	948	957	963	969	993	1006	1010	1014	1018
1	0.01	4052	4999	5404	5624	5764	5859	5928	5981	6022	6056	6209	6286	6313	6340	6366
1	0.005	16212	19997	21614	22501	23056	23440	23715	23924	24091	24222	24837	25146	25254	25358	25466
1	0.001	405312	499725	540257	562668	576496	586033	593185	597954	602245	605583	620842	628471	631332	634193	636578
2	0.05	18.51	19.00	19.16	19.25	19.30	19.33	19.35	19.37	19.38	19.40	19.45	19.47	19.48	19.49	19.50
2	0.025	38.51	39.00	39.17	39.25	39.30	39.33	39.36	39.37	39.39	39.40	39.45	39.47	39.48	39.49	39.50
2	0.01	98.50	99.00	99.16	99.25	99.30	99.33	99.36	99.38	99.39	99.40	99.45	99.48	99.48	99.49	99.50
2	0.005	198.50	199.01	199.16	199.24	199.30	199.33	199.36	199.38	199.39	199.39	199.45	199.48	199.48	199.49	199.51
2	0.001	998.38	998.84	999.31	999.31	999.31	999.31	999.31	999.31	999.31	999.31	999.31	999.31	999.31	999.31	999.31
3	0.05	10.13	9.55	9.28	9.12	9.01	8.94	8.89	8.85	8.81	8.79	8.66	8.59	8.57	8.55	8.53
3	0.025	17.44	16.04	15.44	15.10	14.88	14.73	14.62	14.54	14.47	14.42	14.17	14.04	13.99	13.95	13.90
3	0.01	34.12	30.82	29.46	28.71	28.24	27.91	27.67	27.49	27.34	27.23	26.69	26.41	26.32	26.22	26.13
3	0.005	55.55	49.80	47.47	46.20	45.39	44.84	44.43	44.13	43.88	43.68	42.78	42.31	42.15	41.99	41.83
3	0.001	167.06	148.49	141.10	137.08	134.58	132.83	131.61	130.62	129.86	129.22	126.43	124.97	124.45	123.98	123.46
4	0.05	7.71	6.94	6.59	6.39	6.26	6.16	6.09	6.04	6.00	5.96	5.80	5.72	5.69	5.66	5.63
4	0.025	12.22	10.65	9.98	9.60	9.36	9.20	9.07	8.98	8.90	8.84	8.56	8.41	8.36	8.31	8.26
4	0.01	21.20	18.00	16.69	15.98	15.52	15.21	14.98	14.80	14.66	14.55	14.02	13.75	13.65	13.56	13.46
4	0.005	31.33	26.28	24.26	23.15	22.46	21.98	21.62	21.35	21.14	20.97	20.17	19.75	19.61	19.47	19.32
4	0.001	74.13	61.25	56.17	53.43	51.72	50.52	49.65	49.00	48.47	48.05	46.10	45.08	44.75	44.40	44.05
5	0.05	6.61	5.79	5.41	5.19	5.05	4.95	4.88	4.82	4.77	4.74	4.56	4.46	4.43	4.40	4.37
5	0.025	10.01	8.43	7.76	7.39	7.15	6.98	6.85	6.76	6.68	6.62	6.33	6.18	6.12	6.07	6.02
5	0.01	16.26	13.27	12.06	11.39	10.97	10.67	10.46	10.29	10.16	10.05	9.55	9.29	9.20	9.11	9.02
5	0.005	22.78	18.31	16.53	15.56	14.94	14.51	14.20	13.96	13.77	13.62	12.90	12.53	12.40	12.27	12.14
5	0.001	47.18	37.12	33.20	31.08	29.75	28.83	28.17	27.65	27.24	26.91	25.39	24.60	24.33	24.06	23.79
6	0.05	5.99	5.14	4.76	4.53	4.39	4.28	4.21	4.15	4.10	4.06	3.87	3.77	3.74	3.70	3.67
6	0.025	8.81	7.26	6.60	6.23	5.99	5.82	5.70	5.60	5.52	5.46	5.17	5.01	4.96	4.90	4.85
6	0.01	13.75	10.92	9.78	9.15	8.75	8.47	8.26	8.10	7.98	7.87	7.40	7.14	7.06	6.97	6.88
6	0.005	18.63	14.54	12.92	12.03	11.46	11.07	10.79	10.57	10.39	10.25	9.59	9.24	9.12	9.00	8.88
6	0.001	35.51	27.00	23.71	21.92	20.80	20.03	19.46	19.03	18.69	18.41	17.12	16.44	16.21	15.98	15.75
7	0.05	5.59	4.74	4.35	4.12	3.97	3.87	3.79	3.73	3.68	3.64	3.44	3.34	3.30	3.27	3.23
7	0.025	8.07	6.54	5.89	5.52	5.29	5.12	4.99	4.90	4.82	4.76	4.47	4.31	4.25	4.20	4.14
7	0.01	12.25	9.55	8.45	7.85	7.46	7.19	6.99	6.84	6.72	6.62	6.16	5.91	5.82	5.74	5.65
7	0.005	16.24	12.40	10.88	10.05	9.52	9.16	8.89	8.68	8.51	8.38	7.75	7.42	7.31	7.19	7.08
7	0.001	29.25	21.69	18.77	17.20	16.21	15.52	15.02	14.63	14.33	14.08	12.93	12.33	12.12	11.91	11.70
8	0.05	5.32	4.46	4.07	3.84	3.69	3.58	3.50	3.44	3.39	3.35	3.15	3.04	3.01	2.97	2.93
8	0.025	7.57	6.06	5.42	5.05	4.82	4.65	4.53	4.43	4.36	4.30	4.00	3.84	3.78	3.73	3.67
8	0.01	11.26	8.65	7.59	7.01	6.63	6.37	6.18	6.03	5.91	5.81	5.36	5.12	5.03	4.95	4.86
8	0.005	14.69	11.04	9.60	8.81	8.30	7.95	7.69	7.50	7.34	7.21	6.61	6.29	6.18	6.06	5.95
8	0.001	25.41	18.49	15.83	14.39	13.48	12.86	12.40	12.05	11.77	11.54	10.48	9.92	9.73	9.53	9.33
9	0.05	5.12	4.26	3.86	3.63	3.48	3.37	3.29	3.23	3.18	3.14	2.94	2.83	2.79	2.75	2.71
9	0.025	7.21	5.71	5.08	4.72	4.48	4.32	4.20	4.10	4.03	3.96	3.67	3.51	3.45	3.39	3.33
9	0.01	10.56	8.02	6.99	6.42	6.06	5.80	5.61	5.47	5.35	5.26	4.81	4.57	4.48	4.40	4.31
9	0.005	13.61	10.11	8.72	7.96	7.47	7.13	6.88	6.69	6.54	6.42	5.83	5.52	5.41	5.30	5.19
9	0.001	22.86	16.39	13.90	12.56	11.71	11.13	10.70	10.37	10.11	9.89	8.90	8.37	8.19	8.00	7.81

10	0.05	4.96	4.10	3.71	3.48	3.33	3.22	3.14	3.07	3.02	2.98	2.77	2.66	2.62	2.58	2.54
10	0.025	6.94	5.46	4.83	4.47	4.24	4.07	3.95	3.85	3.78	3.72	3.42	3.26	3.20	3.14	3.08
10	0.01	10.04	7.56	6.55	5.99	5.64	5.39	5.20	5.06	4.94	4.85	4.41	4.17	4.08	4.00	3.91
10	0.005	12.83	9.43	8.08	7.34	6.87	6.54	6.30	6.12	5.97	5.85	5.27	4.97	4.86	4.75	4.64
10	0.001	21.04	14.90	12.55	11.28	10.48	9.93	9.52	9.20	8.96	8.75	7.80	7.30	7.12	6.94	6.76
12	0.05	4.75	3.89	3.49	3.26	3.11	3.00	2.91	2.85	2.80	2.75	2.54	2.43	2.38	2.34	2.30
12	0.025	6.55	5.10	4.47	4.12	3.89	3.73	3.61	3.51	3.44	3.37	3.07	2.91	2.85	2.79	2.72
12	0.01	9.33	6.93	5.95	5.41	5.06	4.82	4.64	4.50	4.39	4.30	3.86	3.62	3.54	3.45	3.36
12	0.005	11.75	8.51	7.23	6.52	6.07	5.76	5.52	5.35	5.20	5.09	4.53	4.23	4.12	4.01	3.90
12	0.001	18.64	12.97	10.80	9.63	8.89	8.38	8.00	7.71	7.48	7.29	6.40	5.93	5.76	5.59	5.42
14	0.05	4.60	3.74	3.34	3.11	2.96	2.85	2.76	2.70	2.65	2.60	2.39	2.27	2.22	2.18	2.13
14	0.025	6.30	4.86	4.24	3.89	3.66	3.50	3.38	3.29	3.21	3.15	2.84	2.67	2.61	2.55	2.49
14	0.01	8.86	6.51	5.56	5.04	4.69	4.46	4.28	4.14	4.03	3.94	3.51	3.27	3.18	3.09	3.00
14	0.005	11.06	7.92	6.68	6.00	5.56	5.26	5.03	4.86	4.72	4.60	4.06	3.76	3.66	3.55	3.44
14	0.001	17.14	11.78	9.73	8.62	7.92	7.44	7.08	6.80	6.58	6.40	5.56	5.10	4.94	4.77	4.60
16	0.05	4.49	3.63	3.24	3.01	2.85	2.74	2.66	2.59	2.54	2.49	2.28	2.15	2.11	2.06	2.01
16	0.025	6.12	4.69	4.08	3.73	3.50	3.34	3.22	3.12	3.05	2.99	2.68	2.51	2.45	2.38	2.32
16	0.01	8.53	6.23	5.29	4.77	4.44	4.20	4.03	3.89	3.78	3.69	3.26	3.02	2.93	2.84	2.75
16	0.005	10.58	7.51	6.30	5.64	5.21	4.91	4.69	4.52	4.38	4.27	3.73	3.44	3.33	3.22	3.11
16	0.001	16.12	10.97	9.01	7.94	7.27	6.80	6.46	6.20	5.98	5.81	4.99	4.54	4.39	4.23	4.06
18	0.05	4.41	3.55	3.16	2.93	2.77	2.66	2.58	2.51	2.46	2.41	2.19	2.06	2.02	1.97	1.92
18	0.025	5.98	4.56	3.95	3.61	3.38	3.22	3.10	3.01	2.93	2.87	2.56	2.38	2.32	2.26	2.19
18	0.01	8.29	6.01	5.09	4.58	4.25	4.01	3.84	3.71	3.60	3.51	3.08	2.84	2.75	2.66	2.57
18	0.005	10.22	7.21	6.03	5.37	4.96	4.66	4.44	4.28	4.14	4.03	3.50	3.20	3.10	2.99	2.87
18	0.001	15.38	10.39	8.49	7.46	6.81	6.35	6.02	5.76	5.56	5.39	4.59	4.15	4.00	3.84	3.67
20	0.05	4.35	3.49	3.10	2.87	2.71	2.60	2.51	2.45	2.39	2.35	2.12	1.99	1.95	1.90	1.84
20	0.025	5.87	4.46	3.86	3.51	3.29	3.13	3.01	2.91	2.84	2.77	2.46	2.29	2.22	2.16	2.09
20	0.01	8.10	5.85	4.94	4.43	4.10	3.87	3.70	3.56	3.46	3.37	2.94	2.69	2.61	2.52	2.42
20	0.005	9.94	6.99	5.82	5.17	4.76	4.47	4.26	4.09	3.96	3.85	3.32	3.02	2.92	2.81	2.69
20	0.001	14.82	9.95	8.10	7.10	6.46	6.02	5.69	5.44	5.24	5.08	4.29	3.86	3.70	3.54	3.38
25	0.05	4.24	3.39	2.99	2.76	2.60	2.49	2.40	2.34	2.28	2.24	2.01	1.87	1.82	1.77	1.71
25	0.025	5.69	4.29	3.69	3.35	3.13	2.97	2.85	2.75	2.68	2.61	2.30	2.12	2.05	1.98	1.91
25	0.01	7.77	5.57	4.68	4.18	3.85	3.63	3.46	3.32	3.22	3.13	2.70	2.45	2.36	2.27	2.17
25	0.005	9.48	6.60	5.46	4.84	4.43	4.15	3.94	3.78	3.64	3.54	3.01	2.72	2.61	2.50	2.38
25	0.001	13.88	9.22	7.45	6.49	5.89	5.46	5.15	4.91	4.71	4.56	3.79	3.37	3.22	3.06	2.89
30	0.05	4.17	3.32	2.92	2.69	2.53	2.42	2.33	2.27	2.21	2.16	1.93	1.79	1.74	1.68	1.62
30	0.025	5.57	4.18	3.59	3.25	3.03	2.87	2.75	2.65	2.57	2.51	2.20	2.01	1.94	1.87	1.79
30	0.01	7.56	5.39	4.51	4.02	3.70	3.47	3.30	3.17	3.07	2.98	2.55	2.30	2.21	2.11	2.01
30	0.005	9.18	6.35	5.24	4.62	4.23	3.95	3.74	3.58	3.45	3.34	2.82	2.52	2.42	2.30	2.18
30	0.001	13.29	8.77	7.05	6.12	5.53	5.12	4.82	4.58	4.39	4.24	3.49	3.07	2.92	2.76	2.59
40	0.05	4.08	3.23	2.84	2.61	2.45	2.34	2.25	2.18	2.12	2.08	1.84	1.69	1.64	1.58	1.51
40	0.025	5.42	4.05	3.46	3.13	2.90	2.74	2.62	2.53	2.45	2.39	2.07	1.88	1.80	1.72	1.64
40	0.01	7.31	5.18	4.31	3.83	3.51	3.29	3.12	2.99	2.89	2.80	2.37	2.11	2.02	1.92	1.80
40	0.005	8.83	6.07	4.98	4.37	3.99	3.71	3.51	3.35	3.22	3.12	2.60	2.30	2.18	2.06	1.93
40	0.001	12.61	8.25	6.59	5.70	5.13	4.73	4.44	4.21	4.02	3.87	3.15	2.73	2.57	2.41	2.23

60	0.05	4.00	3.15	2.76	2.53	2.37	2.25	2.17	2.10	2.04	1.99	1.75	1.59	1.53	1.47	1.39
60	0.025	5.29	3.93	3.34	3.01	2.79	2.63	2.51	2.41	2.33	2.27	1.94	1.74	1.67	1.58	1.48
60	0.01	7.08	4.98	4.13	3.65	3.34	3.12	2.95	2.82	2.72	2.63	2.20	1.94	1.84	1.73	1.60
60	0.005	8.49	5.79	4.73	4.14	3.76	3.49	3.29	3.13	3.01	2.90	2.39	2.08	1.96	1.83	1.69
60	0.001	11.97	7.77	6.17	5.31	4.76	4.37	4.09	3.86	3.69	3.54	2.83	2.41	2.25	2.08	1.89
120	0.05	3.92	3.07	2.68	2.45	2.29	2.18	2.09	2.02	1.96	1.91	1.66	1.50	1.43	1.35	1.25
120	0.025	5.15	3.80	3.23	2.89	2.67	2.52	2.39	2.30	2.22	2.16	1.82	1.61	1.53	1.43	1.31
120	0.01	6.85	4.79	3.95	3.48	3.17	2.96	2.79	2.66	2.56	2.47	2.03	1.76	1.66	1.53	1.38
120	0.005	8.18	5.54	4.50	3.92	3.55	3.28	3.09	2.93	2.81	2.71	2.19	1.87	1.75	1.61	1.43
120	0.001	11.38	7.32	5.78	4.95	4.42	4.04	3.77	3.55	3.38	3.24	2.53	2.11	1.95	1.77	1.54
vc	0.05	3.84	3.00	2.60	2.37	2.21	2.10	2.01	1.94	1.88	1.83	1.57	1.39	1.32	1.22	1.00
vc	0.025	5.02	3.69	3.12	2.79	2.57	2.41	2.29	2.19	2.11	2.05	1.71	1.48	1.39	1.27	1.00
vc	0.01	6.63	4.61	3.78	3.32	3.02	2.80	2.64	2.51	2.41	2.32	1.88	1.59	1.47	1.32	1.00
vc	0.005	7.88	5.30	4.28	3.72	3.35	3.09	2.90	2.74	2.62	2.52	2.00	1.67	1.53	1.36	1.01
vc	0.001	10.83	6.91	5.42	4.62	4.10	3.74	3.47	3.27	3.10	2.96	2.27	1.84	1.66	1.45	1.01

Bảng A5 Điểm phần trăm của phân phối χ^2

Cải biên từ Bảng 8 của White et al. (1979) có xin phép tác giả và nhà xuất bản.
d.f.=1 khía so sánh hai tỉ lệ (kiểm định χ^2 2×2 hay χ^2 Mantel Haenzel) hay trong đánh giá
khuynh hướng, điểm phần trăm của kiểm định 2 đuôi. Có thể làm kiểm định một đuôi bằng cách
chia đôi giá trị P (không thể dùng khái niệm một hoặc hai đuôi cho độ tự do lớn hơn)

d.f.	Giá trị P							
	0.5	0.25	0.1	0.05	0.025	0.01	0.005	0.001
1	0.45	1.32	2.71	3.84	5.02	6.63	7.88	10.83
2	1.39	2.77	4.61	5.99	7.38	9.21	10.60	13.82
3	2.37	4.11	6.25	7.81	9.35	11.34	12.84	16.27
4	3.36	5.39	7.78	9.49	11.14	13.28	14.86	18.47
5	4.35	6.63	9.24	11.07	12.83	15.09	16.75	20.51
6	5.35	7.84	10.64	12.59	14.45	16.81	18.55	22.46
7	6.35	9.04	12.02	14.07	16.01	18.48	20.28	24.32
8	7.34	10.22	13.36	15.51	17.53	20.09	21.95	26.12
9	8.34	11.39	14.68	16.92	19.02	21.67	23.59	27.88
10	9.34	12.55	15.99	18.31	20.48	23.21	25.19	29.59
11	10.34	13.70	17.28	19.68	21.92	24.73	26.76	31.26
12	11.34	14.85	18.55	21.03	23.34	26.22	28.30	32.91
13	12.34	15.98	19.81	22.36	24.74	27.69	29.82	34.53
14	13.34	17.12	21.06	23.68	26.12	29.14	31.32	36.12
15	14.34	18.25	22.31	25.00	27.49	30.58	32.80	37.70
16	15.34	19.37	23.54	26.30	28.85	32.00	34.27	39.25
17	16.34	20.49	24.77	27.59	30.19	33.41	35.72	40.79
18	17.34	21.60	25.99	28.87	31.53	34.81	37.16	42.31
19	18.34	22.72	27.20	30.14	32.85	36.19	38.58	43.82
20	19.34	23.83	28.41	31.41	34.17	37.57	40.00	45.31
21	20.34	24.93	29.62	32.67	35.48	38.93	41.40	46.80
22	21.34	26.04	30.81	33.92	36.78	40.29	42.80	48.27
23	22.34	27.14	32.01	35.17	38.08	41.64	44.18	49.73
24	23.34	28.24	33.20	36.42	39.36	42.98	45.56	51.18
25	24.34	29.34	34.38	37.65	40.65	44.31	46.93	52.62
26	25.34	30.43	35.56	38.89	41.92	45.64	48.29	54.05
27	26.34	31.53	36.74	40.11	43.19	46.96	49.65	55.48
28	27.34	32.62	37.92	41.34	44.46	48.28	50.99	56.89
29	28.34	33.71	39.09	42.56	45.72	49.59	52.34	58.30
30	29.34	34.80	40.26	43.77	46.98	50.89	53.67	59.70
40	39.34	45.62	51.81	55.76	59.34	63.69	66.77	73.40
50	49.33	56.33	63.17	67.50	71.42	76.15	79.49	86.66
60	59.33	66.98	74.40	79.08	83.30	88.38	91.95	99.61
70	69.33	77.58	85.53	90.53	95.02	100.43	104.21	112.32
80	79.33	88.13	96.58	101.88	106.63	112.33	116.32	124.84
90	89.33	98.65	107.57	113.15	118.14	124.12	128.30	137.21
100	99.33	109.14	118.50	124.34	129.56	135.81	140.17	149.45

Bảng A6 Probits

Cải biên từ Bảng A4 của Pearson và Hartley (1966) có xin phép Biometrika Trustees
Probits = giá trị của phân phối bình thường tương ứng với phần trăm tích lũy

%	phần mười của %									
	0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
0	-	-3.09	-2.88	-2.75	-2.65	-2.58	-2.51	-2.46	-2.41	-2.37
1	-2.33	-2.29	-2.26	-2.23	-2.20	-2.17	-2.14	-2.12	-2.10	-2.07
2	-2.05	-2.03	-2.01	-2.00	-1.98	-1.96	-1.94	-1.93	-1.91	-1.90
3	-1.88	-1.87	-1.85	-1.84	-1.83	-1.81	-1.80	-1.79	-1.77	-1.76
4	-1.75	-1.74	-1.73	-1.72	-1.71	-1.70	-1.68	-1.67	-1.66	-1.65
5	-1.64	-1.64	-1.63	-1.62	-1.61	-1.60	-1.59	-1.58	-1.57	-1.56
6	-1.55	-1.55	-1.54	-1.53	-1.52	-1.51	-1.51	-1.50	-1.49	-1.48
7	-1.48	-1.47	-1.46	-1.45	-1.45	-1.44	-1.43	-1.43	-1.42	-1.41
8	-1.41	-1.40	-1.39	-1.39	-1.38	-1.37	-1.37	-1.36	-1.35	-1.35
9	-1.34	-1.33	-1.33	-1.32	-1.32	-1.31	-1.30	-1.30	-1.29	-1.29
10	-1.28	-1.28	-1.27	-1.26	-1.26	-1.25	-1.25	-1.24	-1.24	-1.23
11	-1.23	-1.22	-1.22	-1.21	-1.21	-1.20	-1.20	-1.19	-1.19	-1.18
12	-1.17	-1.17	-1.17	-1.16	-1.16	-1.15	-1.15	-1.14	-1.14	-1.13
13	-1.13	-1.12	-1.12	-1.11	-1.11	-1.10	-1.10	-1.09	-1.09	-1.08
14	-1.08	-1.08	-1.07	-1.07	-1.06	-1.06	-1.05	-1.05	-1.05	-1.04
15	-1.04	-1.03	-1.03	-1.02	-1.02	-1.02	-1.01	-1.01	-1.00	-1.00
16	-0.99	-0.99	-0.99	-0.98	-0.98	-0.97	-0.97	-0.97	-0.96	-0.96
17	-0.95	-0.95	-0.95	-0.94	-0.94	-0.93	-0.93	-0.93	-0.92	-0.92
18	-0.92	-0.91	-0.91	-0.90	-0.90	-0.90	-0.89	-0.89	-0.89	-0.88
19	-0.88	-0.87	-0.87	-0.87	-0.86	-0.86	-0.86	-0.85	-0.85	-0.85
20	-0.84	-0.84	-0.83	-0.83	-0.83	-0.82	-0.82	-0.82	-0.81	-0.81
21	-0.81	-0.80	-0.80	-0.80	-0.79	-0.79	-0.79	-0.78	-0.78	-0.78
22	-0.77	-0.77	-0.77	-0.76	-0.76	-0.76	-0.75	-0.75	-0.75	-0.74
23	-0.74	-0.74	-0.73	-0.73	-0.73	-0.72	-0.72	-0.72	-0.71	-0.71
24	-0.71	-0.70	-0.70	-0.70	-0.69	-0.69	-0.69	-0.68	-0.68	-0.68
25	-0.67	-0.67	-0.67	-0.67	-0.66	-0.66	-0.66	-0.65	-0.65	-0.65
26	-0.64	-0.64	-0.64	-0.63	-0.63	-0.63	-0.62	-0.62	-0.62	-0.62
27	-0.61	-0.61	-0.61	-0.60	-0.60	-0.60	-0.59	-0.59	-0.59	-0.59
28	-0.58	-0.58	-0.58	-0.57	-0.57	-0.57	-0.57	-0.56	-0.56	-0.56
29	-0.55	-0.55	-0.55	-0.54	-0.54	-0.54	-0.54	-0.53	-0.53	-0.53
30	-0.52	-0.52	-0.52	-0.52	-0.51	-0.51	-0.51	-0.50	-0.50	-0.50
31	-0.50	-0.49	-0.49	-0.49	-0.48	-0.48	-0.48	-0.48	-0.47	-0.47
32	-0.47	-0.46	-0.46	-0.46	-0.46	-0.45	-0.45	-0.45	-0.45	-0.44
33	-0.44	-0.44	-0.43	-0.43	-0.43	-0.43	-0.42	-0.42	-0.42	-0.42
34	-0.41	-0.41	-0.41	-0.40	-0.40	-0.40	-0.40	-0.39	-0.39	-0.39
35	-0.39	-0.38	-0.38	-0.38	-0.37	-0.37	-0.37	-0.37	-0.36	-0.36
36	-0.36	-0.36	-0.35	-0.35	-0.35	-0.35	-0.34	-0.34	-0.34	-0.33
37	-0.33	-0.33	-0.33	-0.32	-0.32	-0.32	-0.32	-0.31	-0.31	-0.31
38	-0.31	-0.30	-0.30	-0.30	-0.29	-0.29	-0.29	-0.29	-0.28	-0.28
39	-0.28	-0.28	-0.27	-0.27	-0.27	-0.27	-0.26	-0.26	-0.26	-0.26
40	-0.25	-0.25	-0.25	-0.25	-0.24	-0.24	-0.24	-0.24	-0.23	-0.23
41	-0.23	-0.22	-0.22	-0.22	-0.22	-0.21	-0.21	-0.21	-0.21	-0.20
42	-0.20	-0.20	-0.20	-0.19	-0.19	-0.19	-0.19	-0.18	-0.18	-0.18
43	-0.18	-0.17	-0.17	-0.17	-0.17	-0.16	-0.16	-0.16	-0.16	-0.15
44	-0.15	-0.15	-0.15	-0.14	-0.14	-0.14	-0.14	-0.13	-0.13	-0.13
45	-0.13	-0.12	-0.12	-0.12	-0.12	-0.11	-0.11	-0.11	-0.11	-0.10
46	-0.10	-0.10	-0.10	-0.09	-0.09	-0.09	-0.09	-0.08	-0.08	-0.08
47	-0.08	-0.07	-0.07	-0.07	-0.07	-0.06	-0.06	-0.06	-0.06	-0.05
48	-0.05	-0.05	-0.05	-0.04	-0.04	-0.04	-0.04	-0.03	-0.03	-0.03
49	-0.03	-0.02	-0.02	-0.02	-0.02	-0.01	-0.01	-0.01	-0.01	0.00
50	0.00	0.00	0.01	0.01	0.01	0.01	0.02	0.02	0.02	0.02

51	0.03	0.03	0.03	0.03	0.04	0.04	0.04	0.04	0.05	0.05
52	0.05	0.05	0.06	0.06	0.06	0.06	0.07	0.07	0.07	0.07
53	0.08	0.08	0.08	0.08	0.09	0.09	0.09	0.09	0.10	0.10
54	0.10	0.10	0.11	0.11	0.11	0.11	0.12	0.12	0.12	0.12
55	0.13	0.13	0.13	0.13	0.14	0.14	0.14	0.14	0.15	0.15
56	0.15	0.15	0.16	0.16	0.16	0.16	0.17	0.17	0.17	0.17
57	0.18	0.18	0.18	0.18	0.19	0.19	0.19	0.19	0.20	0.20
58	0.20	0.20	0.21	0.21	0.21	0.21	0.22	0.22	0.22	0.22
59	0.23	0.23	0.23	0.24	0.24	0.24	0.24	0.25	0.25	0.25
60	0.25	0.26	0.26	0.26	0.26	0.27	0.27	0.27	0.27	0.28
61	0.28	0.28	0.28	0.29	0.29	0.29	0.29	0.30	0.30	0.30
62	0.31	0.31	0.31	0.31	0.32	0.32	0.32	0.32	0.33	0.33
63	0.33	0.33	0.34	0.34	0.34	0.35	0.35	0.35	0.35	0.36
64	0.36	0.36	0.36	0.37	0.37	0.37	0.37	0.38	0.38	0.38
65	0.39	0.39	0.39	0.39	0.40	0.40	0.40	0.40	0.41	0.41
66	0.41	0.42	0.42	0.42	0.42	0.43	0.43	0.43	0.43	0.44
67	0.44	0.44	0.45	0.45	0.45	0.45	0.46	0.46	0.46	0.46
68	0.47	0.47	0.47	0.48	0.48	0.48	0.48	0.49	0.49	0.49
69	0.50	0.50	0.50	0.50	0.51	0.51	0.51	0.52	0.52	0.52
70	0.52	0.53	0.53	0.53	0.54	0.54	0.54	0.54	0.55	0.55
71	0.55	0.56	0.56	0.56	0.57	0.57	0.57	0.57	0.58	0.58
72	0.58	0.59	0.59	0.59	0.59	0.60	0.60	0.60	0.61	0.61
73	0.61	0.62	0.62	0.62	0.62	0.63	0.63	0.63	0.64	0.64
74	0.64	0.65	0.65	0.65	0.66	0.66	0.66	0.67	0.67	0.67
75	0.67	0.68	0.68	0.68	0.69	0.69	0.69	0.70	0.70	0.70
76	0.71	0.71	0.71	0.72	0.72	0.72	0.73	0.73	0.73	0.74
77	0.74	0.74	0.75	0.75	0.75	0.76	0.76	0.76	0.77	0.77
78	0.77	0.78	0.78	0.78	0.79	0.79	0.79	0.80	0.80	0.80
79	0.81	0.81	0.81	0.82	0.82	0.82	0.83	0.83	0.83	0.84
80	0.84	0.85	0.85	0.85	0.86	0.86	0.86	0.87	0.87	0.87
81	0.88	0.88	0.89	0.89	0.89	0.90	0.90	0.90	0.91	0.91
82	0.92	0.92	0.92	0.93	0.93	0.93	0.94	0.94	0.95	0.95
83	0.95	0.96	0.96	0.97	0.97	0.97	0.98	0.98	0.99	0.99
84	0.99	1.00	1.00	1.01	1.01	1.02	1.02	1.02	1.03	1.03
85	1.04	1.04	1.05	1.05	1.05	1.06	1.06	1.07	1.07	1.08
86	1.08	1.08	1.09	1.09	1.10	1.10	1.11	1.11	1.12	1.12
87	1.13	1.13	1.14	1.14	1.15	1.15	1.16	1.16	1.17	1.17
88	1.17	1.18	1.19	1.19	1.20	1.20	1.21	1.21	1.22	1.22
89	1.23	1.23	1.24	1.24	1.25	1.25	1.26	1.26	1.27	1.28
90	1.28	1.29	1.29	1.30	1.30	1.31	1.32	1.32	1.33	1.33
91	1.34	1.35	1.35	1.36	1.37	1.37	1.38	1.39	1.39	1.40
92	1.41	1.41	1.42	1.43	1.43	1.44	1.45	1.45	1.46	1.47
93	1.48	1.48	1.49	1.50	1.51	1.51	1.52	1.53	1.54	1.55
94	1.55	1.56	1.57	1.58	1.59	1.60	1.61	1.62	1.63	1.64
95	1.64	1.65	1.66	1.67	1.68	1.70	1.71	1.72	1.73	1.74
96	1.75	1.76	1.77	1.79	1.80	1.81	1.83	1.84	1.85	1.87
97	1.88	1.90	1.91	1.93	1.94	1.96	1.98	2.00	2.01	2.03
98	2.05	2.07	2.10	2.12	2.14	2.17	2.20	2.23	2.26	2.29
99	2.33	2.37	2.41	2.46	2.51	2.58	2.65	2.75	2.88	3.09

Bảng A7 Giá trị tới hạn của kiểm định sắp hạng có dấu cặp đôi Wilcoxon

Tạo lại từ bảng 21 của White et al (1979) có xin phép tác giả và nhà xuất bản

N = số hiệu số khác không

T = số nhỏ nhất trong T+ và T-

Có ý nghĩa nếu $T < \text{giá trị tới hạn}$

N	Giá trị P một bên			
	0.05	0.025	0.01	0.005
	Giá trị P hai bên			
N	0.1	0.05	0.02	0.01
5	1			
6	2	1		
7	4	2	0	
8	6	4	2	0
9	8	6	3	2
10	11	8	5	3
11	14	11	7	5
12	17	14	10	7
13	21	17	13	10
14	26	21	16	13
15	30	25	20	16
16	36	30	24	19
17	41	35	28	23
18	47	40	33	28
19	54	46	38	32
20	60	52	43	37
21	68	59	49	43
22	75	66	56	49
23	83	73	62	55
24	92	81	69	61
25	101	90	77	68
26	110	98	85	76
27	120	107	93	84
28	130	117	102	92
29	141	127	111	100
30	152	137	120	109
31	163	148	130	118
32	175	159	141	128
33	188	171	151	138
34	201	183	162	149
35	214	195	174	160
36	228	208	186	171
37	242	222	198	183
38	256	235	211	195
39	271	250	224	208
40	287	264	238	221
41	303	279	252	234
42	319	295	267	248
43	336	311	281	262
44	353	327	297	277
45	371	344	313	292
46	389	361	329	307
47	408	397	345	323
48	427	397	362	339
49	446	415	380	356
50	466	434	398	373

Bảng A8 Giá trị tới hạn của kiểm định tổng sắp hạng Wilcoxon

Tạo lại từ bảng A7 của Cotton (1974) có xin phép tác giả và nhà xuất bản

n_1, n_2 = cỡ mẫu của hai nhóm

T = tổng các hạng trong nhóm có cỡ mẫu nhỏ

Có ý nghĩa nếu T ở ngoài khoảng tới hạn

n ₁ , n ₂		Giá trị P một bên					
		0.025		0.005		0.0005	
				Giá trị P hai bên			
				0.01		0.001	
2	8	3	19				
2	9	3	21				
2	10	3	23				
2	11	4	24				
2	12	4	26				
2	13	4	28				
2	14	4	30				
2	15	4	32				
2	16	4	34				
2	17	5	35				
2	18	5	37				
2	19	5	39	3	41		
2	20	5	41	3	43		
2	21	6	42	3	45		
2	22	6	44	3	47		
2	23	6	46	3	49		
2	24	6	48	3	51		
2	25	6	50	3	53		
3	5	6	21				
3	6	7	23				
3	7	7	26				
3	8	8	28				
3	9	8	31	6	33		
3	10	9	33	6	36		
3	11	9	36	6	39		
3	12	10	38	7	41		
3	13	10	41	7	44		
3	14	11	43	7	47		
3	15	11	46	8	49		
3	16	12	48	8	52		
3	17	12	51	8	55		
3	18	13	53	8	58		
3	19	13	56	9	60		
3	20	14	58	9	63		
3	21	14	61	9	66	6	69
3	22	15	63	10	68	6	72
3	23	15	66	10	71	6	75
3	24	16	68	10	74	6	78
3	25	16	71	11	76	6	81

n ₁ , n ₂		0.05		0.01		0.001	
4	4	10	26				
4	5	11	29				
4	6	12	32	10	34		
4	7	13	35	10	38		
4	8	14	38	11	41		
4	9	15	41	11	45		
4	10	15	45	12	48		
4	11	16	48	12	52		
4	12	17	51	13	55		
4	13	18	54	14	58	10	62
4	14	19	57	14	62	10	66
4	15	20	60	15	65	10	70
4	16	21	63	15	69	11	73
4	17	21	67	16	72	11	77
4	18	22	70	16	76	11	81
4	19	23	73	17	79	12	84
4	20	24	76	18	82	12	88
4	21	25	79	18	86	12	92
4	22	26	82	19	89	13	95
4	23	27	85	19	93	13	99
4	24	28	88	20	96	13	103
4	25	28	92	20	100	14	106
5	5	17	38	15	40		
5	6	18	42	16	44		
5	7	20	45	17	48		
5	8	21	49	17	53		
5	9	22	53	18	57	15	60
5	10	23	57	19	61	15	65
5	11	24	61	20	65	16	69
5	12	26	64	21	69	16	74
5	13	27	68	22	73	17	78
5	14	28	72	22	78	17	83
5	15	29	76	23	82	18	87
5	16	31	79	24	86	18	92
5	17	32	83	25	90	19	96
5	18	33	87	26	94	19	101
5	19	34	91	27	98	20	105
5	20	35	95	28	102	20	110
5	21	37	98	29	106	21	114
5	22	38	102	29	111	21	119
5	23	39	106	30	115	22	123
5	24	40	110	31	119	23	127
5	25	42	103	32	123	23	132

n_1, n_2		0.05		0.01		0.001	
6	6	26	52	23	55		
6	7	27	57	24	60		
6	8	29	61	25	65	21	69
6	9	31	65	26	70	22	74
6	10	32	70	27	75	23	79
6	11	34	74	28	80	23	85
6	12	35	79	30	84	24	90
6	13	37	83	31	89	25	95
6	14	38	88	32	94	26	100
6	15	40	92	33	99	26	106
6	16	42	96	34	104	27	111
6	17	43	101	36	108	28	116
6	18	45	105	37	113	29	121
6	19	46	110	38	118	29	127
6	20	48	114	39	123	30	132
6	21	50	118	40	128	31	137
6	22	51	123	42	132	32	142
6	23	53	127	43	137	33	147
6	24	55	131	44	142	34	152
7	7	36	69	32	73	28	77
7	8	38	74	34	78	29	83
7	9	40	79	35	84	30	89
7	10	42	84	37	89	31	95
7	11	44	89	38	95	32	101
7	12	46	94	40	100	33	107
7	13	48	99	41	106	34	113
7	14	50	104	43	111	35	119
7	15	52	109	44	117	36	125
7	16	54	114	46	122	37	131
7	17	56	119	47	128	38	137
7	18	58	124	49	133	39	143
7	19	60	129	50	139	41	148
7	20	62	134	52	144	42	154
7	21	64	139	53	150	43	160
7	22	66	144	55	155	44	166
7	23	68	149	57	160	45	172
8	8	49	87	43	93	38	98
8	9	51	93	45	99	40	104
8	10	53	99	47	105	41	111
8	11	55	105	49	111	42	118
8	12	58	110	51	117	43	125
8	13	60	116	53	123	45	131
8	14	63	121	54	130	46	138
8	15	65	127	56	136	47	145
8	16	67	133	58	142	49	151
8	17	70	138	60	148	50	158
8	18	72	144	62	154	51	165
8	19	74	150	64	160	53	171
8	20	77	155	66	166	54	178
8	21	79	161	68	172	56	184
8	22	82	166	70	178	57	191

n_1, n_2		0.05		0.01		0.001	
9	9	63	108	56	115	50	121
9	10	65	115	58	122	52	128
9	11	68	121	61	128	53	136
9	12	71	127	63	135	55	143
9	13	73	134	65	142	56	151
9	14	76	140	67	149	58	158
9	15	79	146	70	155	60	165
9	16	82	152	72	162	61	173
9	17	84	159	74	169	63	180
9	18	87	165	76	176	65	187
9	19	90	171	78	183	66	195
9	20	93	177	81	189	68	202
9	21	95	184	83	196	70	209
10	10	78	132	71	139	63	147
10	11	81	139	74	146	65	155
10	12	85	145	76	154	67	163
10	13	88	152	79	161	69	171
10	14	91	159	81	169	71	179
10	15	94	166	84	176	73	187
10	16	97	173	86	184	75	195
10	17	100	180	89	191	77	203
10	18	103	187	92	198	79	211
10	19	107	193	94	206	81	219
10	20	110	200	97	213	83	227
11	11	96	157	87	166	78	175
11	12	99	165	90	174	81	183
11	13	103	172	93	182	83	192
11	14	106	180	96	190	85	201
11	15	110	187	99	198	87	210
11	16	114	194	102	206	90	218
11	17	117	202	105	214	92	227
11	18	121	209	108	222	94	236
11	19	124	217	111	230	97	244
12	12	115	185	106	194	95	205
12	13	119	193	109	203	98	214
12	14	123	201	112	212	100	224
12	15	127	209	115	221	103	233
12	16	131	217	119	229	105	243
12	17	135	225	122	238	108	252
12	18	139	233	125	247	111	261
13	13	137	214	125	226	114	237
13	14	141	223	129	235	116	248
13	15	145	232	133	244	119	258
13	16	150	240	137	253	122	268
13	17	154	249	140	263	125	278
14	14	160	246	147	259	134	272
14	15	164	256	151	269	137	283
14	16	169	265	155	279	140	294
15	15	185	280	171	294	156	309

Bảng A9 Giá trị tới hạn của hệ số tương quan sắp hạng của Spearman

Có ý nghĩa nếu $r_s >$ giá trị tới hạn. Cải biên từ Bảng A8 của Colton (1974) và từ Bảng 17 của White et al. (1979), có xin phép tác giả và nhà xuất bản

số cặp	Giá trị P một bên		Giá trị P một bên	
	0.05	0.01	0.05	0.01
5			0.9	
6	0.086		0.829	0.943
7	0.786		0.714	0.893
8	0.738	0.881	0.643	0.833
9	0.683	0.833	0.6	0.783
10	0.648	0.794	0.564	0.746
12	0.591	0.78	0.506	0.712
14	0.545	0.716	0.456	0.645
16	0.507	0.666	0.425	0.601
18	0.476	0.625	0.399	0.564
20	0.45	0.591	0.377	0.534
22	0.428	0.562	0.359	0.508
24	0.409	0.537	0.343	0.485
26	0.392	0.515	0.329	0.465
28	0.377	0.496	0.317	0.448
30	0.364	0.478	0.306	0.432

Bảng A10 Số ngẫu nhiên

34735	78219	18131	92594	94235	11721	53806	52733	19665	26574	05728	81270	81641
35621	57344	02606	21961	07539	71006	51935	83322	29423	27142	85686	44342	09145
78629	40478	63628	13640	82315	41919	90964	82448	48074	41423	40836	79588	64174
08462	33570	21715	90409	33199	71764	56738	50619	59743	76495	81334	12020	73947
24014	71381	58732	29417	32050	89880	61392	61573	33128	62969	40939	89200	60756
37124	23597	73007	26705	94330	45206	74130	55905	66409	06851	10401	77370	79931
92775	68533	86784	28870	61590	99165	31797	03780	47408	29291	88999	81583	54406
26426	54602	71259	56747	36957	82629	23159	70407	88383	47691	95014	60902	46232
21487	46012	10948	49446	32178	50727	30892	19248	56504	72663	45070	55686	14624
17745	94929	23861	66784	15825	39009	50060	67336	31511	52316	25839	92967	57622
30495	18694	23722	03685	19700	96192	00151	08091	84548	27850	38503	66775	20398
32372	52818	46875	87319	85180	75405	74269	19501	48521	51843	72300	84055	29993
75451	26763	60571	94992	38611	28142	47473	55362	35318	85379	35938	73876	63602
11762	41680	66807	37812	79865	79427	41731	68471	81021	64385	51734	10187	71142
83884	49794	55501	97110	45306	77600	40130	45029	52682	14441	47791	80935	84251
14269	65406	78705	90654	43192	14170	85702	09650	02333	54434	27017	56578	75149
36343	48319	66650	94373	62872	93369	06342	13734	06577	61311	66067	65470	06779
15496	11053	73309	95443	76374	91922	87184	64769	20045	00059	09817	41361	33022
85104	28594	44846	60598	60749	44268	80423	58992	86707	39003	74184	84599	76744
79670	63423	00315	14875	11492	85727	52774	15725	61539	26867	88297	89560	91783
01898	64946	58151	25072	60705	10247	61850	54819	03372	27480	02607	31817	04231
44069	23863	80350	41791	51294	85320	50969	88518	65089	74666	22634	16628	26606
63055	89602	08803	81492	62758	76960	50273	52443	65284	26315	65822	66697	61693
83321	19962	73527	34635	78648	28633	32031	01370	76576	08634	16212	44914	31187
77054	69898	92184	66259	60345	77074	68009	45901	15502	11268	87119	23661	95880
69907	87164	66650	93954	44440	69031	77532	68923	16217	66085	50724	82314	30987
90287	37202	91124	50480	94422	48102	32066	47136	53534	51891	48612	45191	47777
87558	35054	94824	89447	82164	60059	61863	79891	34078	46344	66598	81489	30743
01580	47873	85392	57121	13681	84955	71158	70995	72485	54632	67880	73891	05757
04455	74718	28303	90358	27153	95274	14366	13901	50283	70153	95076	66796	25177
52083	34394	08739	49470	66296	14183	20924	25128	04890	84576	47494	66048	84214
46752	74230	42172	94075	34968	45259	17032	93902	37647	89517	70775	03838	18491
56551	15908	58391	82674	70615	37303	81149	91201	08470	30825	54441	94465	63508
12880	25154	61576	61751	81850	76738	73447	45267	97964	33313	89353	60036	36412
32612	04635	48102	38869	56136	42105	43699	90397	11968	28889	42463	58364	19592
67080	19279	63450	66826	19313	03230	48711	16457	53530	40581	67423	82437	41400
41036	43809	71873	87841	16296	65631	86429	42280	62538	12616	23129	79515	43289
10638	01419	64349	00924	03682	91112	32391	42009	95040	70317	08829	39081	37986
12438	74937	79849	25213	44568	30721	35220	23673	21105	81979	42688	68270	07955
99112	08402	46319	23263	95208	58936	40786	89801	65878	50334	58243	34281	03358
16311	16122	61899	33385	97310	04988	14475	85420	39643	04048	82730	13130	73452
61450	25720	22871	42940	39895	36887	89416	82663	56490	03027	51598	65146	22541
17156	16421	93043	04299	46784	89396	86661	23390	39388	04284	00321	81178	79590
51574	31717	16199	20938	53318	64235	76469	28088	77810	84225	97020	55951	25461
00040	16495	22530	90427	09475	79693	92178	10145	64438	79172	76021	26669	44039
71382	68440	73671	50695	97534	03954	58349	66779	64465	35220	23691	95815	27666
39894	24705	50512	00919	50389	00004	68679	22953	50956	35320	07350	07762	54467
63465	11953	11635	90674	32865	88904	53019	22828	18637	47228	48176	65884	68189
55528	68750	17124	01659	23255	13164	09310	56236	19990	13958	03747	26945	55357
47926	07712	64115	98554	23254	07362	26137	19452	64289	83646	13737	34698	96796
09070	02968	79123	89790	54631	61425	77161	28785	23730	59257	86048	99683	11718
74583	05284	21832	36901	28669	83265	74128	37224	07310	11162	50192	60078	51340
21356	65134	10766	55847	97110	40608	38745	42250	61498	32780	62064	57535	78737
94569	83865	66237	39173	48434	98112	80319	40587	33931	98799	29927	42051	13332
59579	49150	26240	24307	46740	63926	41405	59985	40177	63084	75229	47222	90182

