

คู่มือการคำนวณขนาดตัวอย่างเบื้องต้น

โดย

โดยนางสาววรารักษ์ เนตรพราว

นักจัดการงานทั่วไปปฏิบัติการ กลุ่มงานส่งเสริมการวิจัย

โรงพยาบาลเจริญกรุงประชารักษ์

คำนำ

ด้วยภารกิจของโรงพยาบาลเจริญกรุงประชารักษ์ในการพัฒนาองค์ความรู้และนวัตกรรมทาง การแพทย์ผ่านงานวิจัย ทำให้บุคลากรทางการแพทย์จำเป็นต้องมีทักษะในการทำวิจัยอย่างมีคุณภาพและเป็น ระบบ หนึ่งในขั้นตอนที่สำคัญที่สุดของการวางแผนงานวิจัยคือ การกำหนดขนาดตัวอย่างที่เหมาะสม เนื่องจากจำนวนตัวอย่างที่เพียงพอจะช่วยให้ผลการศึกษาที่ได้มีความน่าเชื่อถือ สามารถเป็นตัวแทนของ ประชากรได้อย่างมีประสิทธิภาพ และมีอำนาจการทดสอบทางสถิติที่เพียงพอต่อการสรุปผลการวิจัย อย่างไรก็ดีตาม บุคลากรจำนวนไม่น้อยประสบปัญหาในการคำนวณขนาดตัวอย่างเนื่องจากข้อจำกัดด้านเวลา และความซับซ้อนของหลักการทางสถิติ คู่มือ "การคำนวณขนาดตัวอย่างเบื้องต้น" ฉบับนี้จึงถูกจัดทำขึ้น โดยกลุ่มงานส่งเสริมการวิจัย เพื่อเป็นเครื่องมือสนับสนุนให้บุคลากรสามารถทำความเข้าใจหลักการพื้นฐาน และดำเนินการคำนวณขนาดตัวอย่างได้ด้วยตนเองอย่างถูกต้อง โดยมุ่งหวังให้คู่มือนี้เป็นแนวทางที่ช่วยลด ขั้นตอนการทำงานวิจัย และส่งเสริมให้บุคลากรในโรงพยาบาลสามารถจัดทำโครงร่างงานวิจัยได้อย่าง รวดเร็วและมีคุณภาพมากยิ่งขึ้น

ผู้เขียนหวังเป็นอย่างยิ่งว่า คู่มือฉบับนี้จะเป็นประโยชน์ในการขับเคลื่อนงานวิจัยของโรงพยาบาลให้ ก้าวหน้าต่อไป

วราภรณ์ เนตรพราว

สารบัญ

	หน้า
Sample size for estimating a finite population proportion	1
Sample size for estimating an infinite population proportion	3
Sample size for comparing two independent population means	5
Sample size for comping two population proportions	6
การคำนวณขนาดตัวอย่าง Logistic regression	8

สูตรการคำนวณขนาดตัวอย่าง

การเลือกใช้สูตรการคำนวณขนาดตัวอย่างเป็นขั้นตอนที่สำคัญยิ่งในการวางแผนงานวิจัย โดยมีหลักการพื้นฐานที่ว่า ผู้วิจัยต้องเลือกใช้สูตรตามวัตถุประสงค์หลักของงานวิจัย

โดยในคู่มือฉบับนี้ได้จัดทำสูตรการคำนวณขนาดตัวอย่างให้สอดคล้องกับวัตถุประสงค์หลักที่มักพบบ่อยในงานวิจัย โดยเป็นการสอนควบคู่ไปกับการใช้งานแอปพลิเคชัน n4studies  เป็นแอปพลิเคชันสำหรับใช้คำนวณขนาดตัวอย่าง

1. วัตถุประสงค์หลัก ต้องการประมาณค่า (Estimation) กล่าวคือผู้วิจัยมีวัตถุประสงค์หลักต้องการศึกษาความชุกของการเกิดโรคหรือต้องการศึกษาอุบัติการณ์ของการเกิดโรค สูตรที่ใช้ในการคำนวณแบ่งออกเป็น 2 สูตร คือกรณีที่ผู้วิจัยทราบจำนวนประชากร (Finite) และ ผู้วิจัยไม่ทราบจำนวนประชากร (Infinite)

1.1 Sample size for estimating a finite population proportion

Sample size for estimating a finite population proportion

[Help](#)

$$n = \frac{Np(1-p)z_{1-\frac{\alpha}{2}}^2}{d^2(N-1) + p(1-p)z_{1-\frac{\alpha}{2}}^2}$$

Population size (N) =

Proportion (p) =

Error (d) =

*p and d must be a range of 0 to 1.

Alpha (α) = Decimal =

Cluster sampling?

Calculate

Clear

N คือ จำนวนกลุ่มประชากรที่ผู้วิจัยต้องการจะศึกษา

p คือ สัดส่วนของการเกิดโรคที่สนใจ ได้จากการทบทวนวรรณกรรมจากงานวิจัยที่มีลักษณะประชากรคล้ายคลึง

d คือ ความคลาดเคลื่อนที่ยอมรับได้ ซึ่งเป็นค่าที่ผู้วิจัยเป็นผู้กำหนดเอง การเลือกค่า d จึงเป็นสิ่งสำคัญอย่างยิ่งและขึ้นอยู่กับการตัดสินใจของผู้วิจัย หากต้องการความแม่นยำสูง (ค่า d มีค่าน้อย) ก็จะต้องใช้ขนาดตัวอย่างที่ใหญ่ขึ้น ในทางกลับกัน หากยอมรับความ

คลาดเคลื่อนได้มากขึ้น (ค่า d มีค่ามาก) ก็สามารถใช้ขนาดตัวอย่างที่เล็กลงได้ ค่าที่นิยมกำหนดคือ 0.05

ตัวอย่างที่ 1 การใช้งาน สูตร Sample size for estimating a finite population proportion

จากการทบทวนวรรณกรรม ได้ทำการศึกษาผู้ป่วยที่มาคลอดในโรงพยาบาลแห่งหนึ่ง พบความชุกของภาวะซีมเศร้าหลังคลอดร้อยละ 15.71 และผู้วิจัยทราบจำนวนผู้ป่วยที่มีมาคลอดในโรงพยาบาลจากศึกษาย้อนหลังจำนวน 2 ปี มีจำนวน 1,500 ราย

ทำการคำนวณค่าจาก แอปพลิเคชัน n4studies จะได้ดังต่อไปนี้

Sample size for estimating a finite population proportion [Help](#)

$$n = \frac{Np(1-p)z_{1-\frac{\alpha}{2}}^2}{d^2(N-1) + p(1-p)z_{1-\frac{\alpha}{2}}^2}$$

Population size (N) = 1,500

Proportion (p) = 0.1571

Error (d) = 0.05

*p and d must be a range of 0 to 1.

Alpha (α) = 0.05 ◊ Decimal = 2 ◊

Cluster sampling? No ◊

Calculate **Clear**

Output:
 $Z(0.975) = 1.96$
 $n = 179.28$
 Thus, sample size = 180

ดังนั้น การศึกษาในครั้งนี้จะศึกษามารดาหลังคลอดบุตร จำนวน 180 คน

1.2 Sample size for estimating an infinite population proportion

Sample size for estimating an infinite population proportion

[Help](#)

$$n = \frac{z_{1-\frac{\alpha}{2}}^2 p(1-p)}{d^2}$$

Proportion (p) =

Error (d) =

**p and d must be a range of 0 to 1.*

Alpha (α) = 0.05 Decimal = 2

Cluster sampling? No

p คือ สัดส่วนของการเกิดโรคที่สนใจ ได้จากการทบทวนวรรณกรรมจากงานวิจัยที่มีลักษณะประชากรคล้ายคลึง

d คือ ความคลาดเคลื่อนที่ยอมรับได้ ซึ่งเป็นค่าที่ผู้วิจัยเป็นผู้กำหนดเอง การเลือกค่า d จึงเป็นสิ่งสำคัญอย่างยิ่งและขึ้นอยู่กับการตัดสินใจของผู้วิจัย หากต้องการความแม่นยำสูง (ค่า d มีค่าน้อย) ก็จะต้องใช้ขนาดตัวอย่างที่ใหญ่ขึ้น ในทางกลับกัน หากยอมรับความคลาดเคลื่อนได้มากขึ้น (ค่า d มีค่ามาก) ก็สามารถใช้นาขนาดตัวอย่างที่เล็กลงได้ ค่าที่นิยมกำหนดคือ 0.05

ตัวอย่างที่ 2 การใช้งาน สูตร Sample size for estimating an infinite population proportion

จากการทบทวนวรรณกรรม ได้ทำการศึกษาผู้ป่วยที่มาคลอดในโรงพยาบาลแห่งหนึ่ง พบความชุกของภาวะซีมเศร้าหลังคลอดร้อยละ 15.71 และผู้วิจัย**ไม่ทราบ**ข้อมูลจำนวนผู้ป่วยที่มีมาคลอดในโรงพยาบาล

ทำการคำนวณค่าจาก แอปพลิเคชัน n4studies จะได้ดังต่อไปนี้

Sample size for estimating an infinite population proportion [Help](#)

$$n = \frac{z_{1-\frac{\alpha}{2}}^2 p(1-p)}{d^2}$$

Proportion (p) =

Error (d) =

**p and d must be a range of 0 to 1.*

Alpha (α) = Decimal =

Cluster sampling?

Output:
 Z(0.975) = 1.96
 n = 203.48
 Thus, sample size = 204

ดังนั้น การศึกษาในครั้งนี้จะศึกษามารดาหลังคลอดบุตร จำนวน 204 คน

- วัตถุประสงค์หลัก ต้องการเปรียบเทียบ outcome ระหว่างกลุ่ม 2 กลุ่ม ที่เป็นอิสระต่อกัน สูตรที่ใช้ในการคำนวณแบ่งออกเป็น 2 สูตร คือกรณีที่ outcome ที่สนใจเป็น binary outcome และกรณีที่ outcome เป็น continuous outcome

2.1 Sample size for comparing two independent population means

Sample size for comparing two independent population means

[Help](#)

$$n_1 = \frac{(z_{1-\frac{\alpha}{2}} + z_{1-\beta})^2 \left[\sigma_1^2 + \frac{\sigma_2^2}{r} \right]}{\Delta^2}$$

$$r = \frac{n_2}{n_1}, \Delta = \mu_1 - \mu_2$$

Mean in group1 (μ_1) =

Mean in group2 (μ_2) =

SD. in group1 (σ_1) =

SD. in group2 (σ_2) =

Ratio (r) =

Alpha (α) = Decimal =

Beta (β) =

Cluster sampling?

μ_1 คือ ค่าเฉลี่ยของ outcome ที่สนใจในกลุ่มที่ 1 หรือกลุ่ม Treatment

μ_2 คือ ค่าเฉลี่ยของ outcome ที่สนใจในกลุ่มที่ 2 หรือกลุ่ม Control

σ_1 คือ ค่าส่วนเบี่ยงเบนมาตรฐานของ outcome ที่สนใจในกลุ่มที่ 1 หรือกลุ่ม Treatment

σ_2 คือ ค่าส่วนเบี่ยงเบนมาตรฐานของ outcome ที่สนใจในกลุ่มที่ 2 หรือกลุ่ม Control

$r = \frac{n_2}{n_1}$ คืออัตราส่วนระหว่างขนาดตัวอย่างในกลุ่ม 2 ต่อกลุ่ม 1 (Control: Treatment)

ตัวอย่างที่ 3 การใช้งาน สูตร Sample size for comparing two independent population means

ผู้วิจัยต้องการทำการวิจัยเพื่อเปรียบเทียบค่าเฉลี่ยของความดันโลหิตของผู้ป่วยสองกลุ่ม คือ กลุ่มที่ได้รับยาใหม่ (กลุ่มทดลอง) และกลุ่มที่ได้รับยามาตรฐาน (กลุ่มควบคุม) จากการทบทวนวรรณกรรม พบว่า ค่าเฉลี่ยของความดันโลหิตของผู้ป่วยกลุ่มที่ได้รับยาใหม่ (กลุ่มทดลอง) เท่ากับ 110 mmHg ค่าเบี่ยงเบนมาตรฐานของความดันโลหิตของผู้ป่วยกลุ่มที่ได้รับยาใหม่ (กลุ่มทดลอง) เท่ากับ 50 mmHg ค่าเฉลี่ยของความดันโลหิตของผู้ป่วยกลุ่มที่ได้รับยามาตรฐาน (กลุ่มควบคุม) เท่ากับ 130 mmHg ค่าเบี่ยงเบนมาตรฐานของความดันโลหิตของผู้ป่วยกลุ่มที่ได้รับยามาตรฐาน (กลุ่มควบคุม) เท่ากับ 30 mmHg อัตราส่วนระหว่างขนาดตัวอย่าง Control ต่อ Treatment เท่ากับ 1:1

ทำการคำนวณค่าจาก แอปพลิเคชัน n4studies จะได้ดังต่อไปนี้

Mean in group1 (μ_1) = 110

Mean in group2 (μ_2) = 130

SD. in group1 (σ_1) = 50

SD. in group2 (σ_2) = 30

Ratio (r) = 1

Alpha (α) = 0.05 $\diamond$ Decimal = 2 $\diamond$

Beta (β) = 0.2 $\diamond$

Cluster sampling? No $\diamond$

Calculate Clear

Output:

$Z(0.95) = 1.96$, $Z(0.8) = 0.84$,
 $n_1 = 66.64$

Thus, sample size in group1 = 67 and group2
 = 67

ดังนั้น การศึกษาในครั้งนี้จะศึกษาขนาดตัวอย่าง จำนวน 134 คน แบ่งออกเป็นกลุ่มที่ได้รับ
 ยามาตรฐาน (กลุ่มควบคุม) 67 ราย และ กลุ่มที่ได้รับยาใหม่ (กลุ่มทดลอง) 67 ราย

2.2 Sample size for comparing two population proportions

Sample size for comparing two population proportions [Help](#)

$$n_1 = \left[\frac{z_{1-\frac{\alpha}{2}} \sqrt{\bar{p}\bar{q} \left(1 + \frac{1}{r}\right)} + z_{1-\beta} \sqrt{p_1 q_1 + \frac{p_2 q_2}{r}}}{\Delta} \right]^2$$

$\Delta = p_1 - p_2$, $\bar{p} = \frac{p_1 + p_2 r}{1 + r}$, $r = \frac{n_2}{n_1}$
 $q_1 = 1 - p_1$, $q_2 = 1 - p_2$, $\bar{q} = 1 - \bar{p}$

Proportion in group1 (p1) = ?

Proportion in group2 (p2) = ?

*p1 and p2 must be a range of 0 to 1.

Ratio (r) = 1

Alpha (α) = 0.05 $\diamond$ Decimal = 2 $\diamond$

Beta (β) = 0.2 $\diamond$

Continuity correction? No $\diamond$

Cluster sampling? No $\diamond$

p_1 คือ proportions ในกลุ่มที่ 1 หรือกลุ่ม Treatment

p_2 คือ proportions ในกลุ่มที่ 2 หรือกลุ่ม Control

$r = \frac{n_2}{n_1}$ คืออัตราส่วนระหว่างขนาดตัวอย่างในกลุ่ม 2 ต่อกลุ่ม 1 (Control: Treatment)

ตัวอย่างที่ 4 การใช้งาน สูตร Sample size for comparing two population proportions

ผู้วิจัยต้องการเปรียบเทียบอัตราการเกิดภาวะเบาหวานขณะตั้งครรภ์ของสตรีตั้งครรภ์ที่มีโรคอ้วนกับไม่มีโรคอ้วนก่อนการตั้งครรภ์ โดยวัดผลจากสัดส่วนของผู้ป่วยที่หายจากโรค จากการทบทวนวรรณกรรม พบว่า พบว่า ผู้หญิงที่มีภาวะอ้วนก่อนการตั้งครรภ์เกิดภาวะแทรกซ้อนเบาหวานขณะตั้งครรภ์ทั้งหมด 17 คน ใน 272 คน คิดเป็นร้อยละ 6.3 และผู้หญิงที่ไม่มีภาวะอ้วนเกิดภาวะแทรกซ้อนเบาหวานขณะตั้งครรภ์ทั้งหมด 44 คน ใน 1,389 คน คิดเป็นร้อยละ 3.1 อัตราส่วนระหว่างขนาดตัวอย่าง เท่ากับ 1:1

ทำการคำนวณค่าจาก แอปพลิเคชัน n4studies จะได้ดังต่อไปนี้

Proportion in group1 (p1) = 0.063

Proportion in group2 (p2) = 0.031

*p1 and p2 must be a range of 0 to 1.

Ratio (r) = 1

Alpha (α) = 0.05 Decimal = 2

Beta (β) = 0.2

Continuity correction? No

Cluster sampling? No

Calculate Clear

Output:
Z(0.975) = 1.96, Z(0.8) = 0.84,
n1 = 684.69,
Thus, sample size for group1 = 685, and group2 = 685

ดังนั้น การศึกษาในครั้งนี้มีสตรีฝากครรภ์โรงพยาบาล เป็นจำนวน 1,370 ราย โดยแบ่งกลุ่มย่อยเป็นสตรีฝากครรภ์ที่มีภาวะอ้วนก่อนการตั้งครรภ์ 685 ราย และไม่มีภาวะอ้วน 685 ราย

การคำนวณขนาดตัวอย่าง Logistic regression (ตัวแปร outcome เป็น binary)

สูตร $n = (\text{cases} * k) / p$ (k คือจำนวนปัจจัยที่สนใจ , p คือความหุคของการเกิดโรคที่สนใจ)

Peduzzi P et al. A simulation study of the number of events per variable in logistic regression analysis. Clin Epidemiol. 1996 Dec;49(12):1373-9.

กรณีตัวแปรต้นเป็น Numerical (continuous) response:

- 30 per predictor [Cohen&Cohen, 1975]

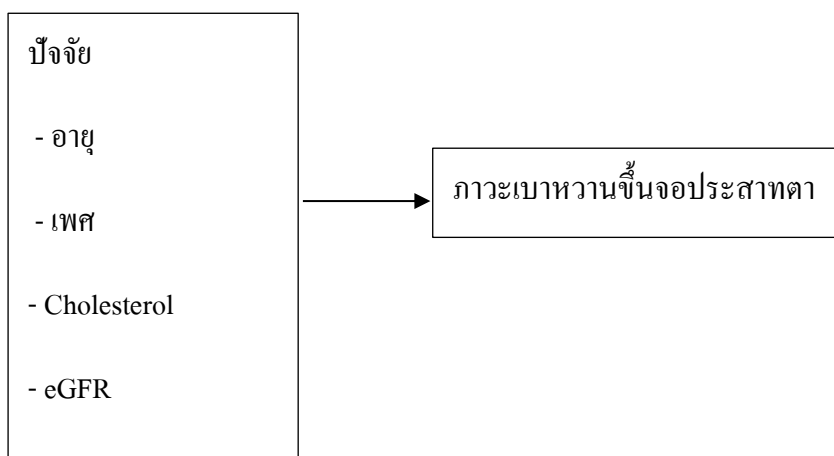
กรณีตัวแปรต้นเป็น Binary response (logistic, Cox's & Poisson regression)

- 10 cases per predictor [Peduzzi et al, 1996]

กรณีตัวแปรต้นเป็น categorical

- $(b-1)*10$ cases per predictor โดยที่ b คือ จำนวนการแบ่งกลุ่มตัวแปร เช่น เพศ แบ่งเป็น 2 กลุ่มชายและหญิง จะใช้ตัวอย่าง $(2-1)*10 = 10$ cases

ตัวอย่างการคำนวณขนาดตัวอย่าง Logistic regression



สูตร $n = (\text{cases} * k) / p$ (k คือจำนวนปัจจัยที่สนใจ , p คือความหุคของการเกิดโรคที่สนใจ)

แทนค่าสูตร $n = (30+10+10+20)/0.40 = 175$ คน

โดยที่ n = ขนาดตัวอย่างในการศึกษานี้

k = จำนวนตัวแปรอิสระ ซึ่งมีจำนวน 4 ตัว

โดยที่ ตัวแปรอายุ เก็บเป็น Numerical ; $n=30$ cases

ตัวแปรเพศ แบ่งเป็น 2 กลุ่ม ชายและหญิง ; $n = (b-1)*10 = (2-1)*10 = 10$ cases

ตัวแปร Cholesterol แบ่งเป็น 2 กลุ่ม well และ poor control

; $n = (b-1)*10 = (2-1)*10 = 10$ cases

ตัวแปร eGFR แบ่งออกเป็น 3 กลุ่ม ได้แก่ mild moderate severe

$$; n = (b-1)*10 = (3-1)*10 = 20 \text{ cases}$$

p = ความชุกของการเกิดภาวะเบาหวานขึ้นจอประสาทตา เท่ากับ 0.40

สูตรการคำนวณขนาดตัวอย่าง เพื่อ Drop out

(จรรยา แก้วกังวาน. (2562). คู่มือนักวิจัยมือใหม่: การประยุกต์ใช้ระบาดวิทยา และชีวสถิติในการวิจัยชีวเวชศาสตร์. กรุงเทพฯ: จี.เอส.เอ็ม.เทรดดิ้ง.)

$$n_{adj} = n / (1 - p)$$

n = ขนาดกลุ่มตัวอย่างที่คำนวณได้

p = สัดส่วนของข้อมูลที่คาดว่าจะสูญหาย

จากการศึกษานี้เพื่อ Drop out 10 % จะได้ขนาดตัวอย่างที่ใช้ในการศึกษานี้เท่ากับ 195 คน

$$n_{adj} = 175 / (1 - 0.1) = 195 \text{ คน}$$