

Basic statistics

Sasivimol Rattanasiri, Ph.D

Clinical Epidemiology Unit

Ramathibodi Hospital

สารบัญ

	หน้า
บทที่ 1 ชนิดของข้อมูล (Types of data)	1
บทที่ 2 การสรุปและนำเสนอข้อมูล (Summarizing data)	3
บทที่ 3 การอ้างอิงทางสถิติ (Statistical inference)	7
บทที่ 4 การทดสอบสมมติฐานสำหรับข้อมูลต่อเนื่อง (Hypothesis testing for continuous data)	12
บทที่ 5 การทดสอบสมมติฐานสำหรับข้อมูลกลุ่ม (Hypothesis testing for categorical data)	23
บทที่ 6 การนำเสนอผลการวิจัย (Reporting results)	28

บทที่ 1

ชนิดของข้อมูล (Types of data)

ในการวิเคราะห์ข้อมูล สิ่งสำคัญที่ผู้วิจัยควรทำความเข้าใจ คือ ชนิดของข้อมูลที่จะนำมาวิเคราะห์ เนื่องจากชนิดของข้อมูลเป็นตัวกำหนดสถิติที่ใช้ในการทดสอบสมมติฐานในงานวิจัย ชนิดของข้อมูลสามารถแบ่งเป็น 2 ประเภทหลักๆ คือ ข้อมูลเชิงคุณภาพ และข้อมูลเชิงปริมาณ

1. ข้อมูลเชิงคุณภาพ หรือข้อมูลกลุ่ม ประกอบด้วย

1.1 Nominal data หมายถึงข้อมูลที่บ่งบอกคุณสมบัติและลักษณะของตัวอย่าง โดยไม่มีการเรียงลำดับความแตกต่างในเชิงปริมาณ เช่น

ข้อมูล	ลักษณะของตัวอย่าง
เพศ	ชาย / หญิง
สถานภาพสมรส	โสด / คู่ / หม้าย / หย่า / แยก
ศาสนา	พุทธ / คริสต์ / อิสลาม / อื่นๆ
กรุ๊ปเลือด	A / B / AB / O

1.2 Ordinal data หมายถึงข้อมูลที่บ่งบอกคุณสมบัติและลักษณะของตัวอย่าง และมีการเรียงลำดับความแตกต่างในเชิงปริมาณ เช่น

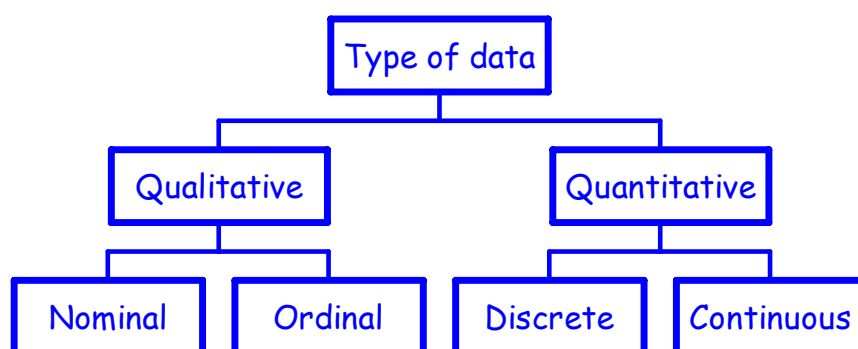
ข้อมูล	ลักษณะของตัวอย่าง
ความเจ็บปวด	ปวดน้อย / ปวดปานกลาง / ปวดมาก
ระดับการศึกษา	ประถมศึกษา / มัธยมศึกษา / มัธยมศึกษาตอนปลาย / ปริญญาตรี
พฤติกรรมการดื่มสุรา	ไม่ดื่ม / ดื่มเป็นบางโอกาส / ดื่ม
ระยะของมะเร็งเต้านม	ระยะที่ 1 / ระยะที่ 2 / ระยะที่ 3 / ระยะที่ 4

ข้อมูลประเภทนี้มีระดับการวัดที่ดีกว่ากลุ่มแรก เนื่องจากมีการจัดเรียงลำดับของแต่ละกลุ่มว่ามีลักษณะมากหรือน้อยกว่ากัน อย่างไรก็ตามข้อมูลชนิดนี้ไม่สามารถนำมาคำนวณหาขนาดของความแตกต่างระหว่างกลุ่มได้

2. ข้อมูลเชิงปริมาณ หรือข้อมูลต่อเนื่อง ประกอบด้วย

2.1 Discrete numerical data หมายถึง ข้อมูลตัวเลขที่มีลักษณะเป็นช่วงๆ ส่วนใหญ่จะเป็นข้อมูลที่ได้จากการนับจำนวนเหตุการณ์ต่างๆ เช่น คะแนนความเจ็บปวด (pain score) จำนวนครั้งที่ตั้งครรภ์ หรือ ระยะเวลาในการฟอกเลือดของผู้ป่วยโรคไตที่ได้รับการรักษาที่ รพ.รามาริบัติ ซึ่งข้อมูลชนิดนี้สามารถจัดอยู่ในกลุ่มข้อมูลเชิงคุณภาพหรือข้อมูลเชิงปริมาณก็ได้ ขึ้นอยู่กับความแตกต่างของตัวเลขในข้อมูลนั้นๆ โดยข้อมูลที่มีความแตกต่างกันมากๆ เช่น ข้อมูลเกี่ยวกับระยะเวลาในการฟอกเลือดของผู้ป่วยโรคไตที่ได้รับการรักษาที่ รพ.รามาริบัติ ซึ่งค่าที่เป็นไปได้มีตั้งแต่เจ็ดเดือน จนถึงสิบปี จึงควรจัดอยู่ในกลุ่มข้อมูลเชิงปริมาณ และใช้หลักการวิเคราะห์ข้อมูลแบบข้อมูลเชิงปริมาณ ในทางกลับกัน ข้อมูลที่มีความแตกต่างกันน้อยๆ เช่น ข้อมูลเกี่ยวกับจำนวนครั้งที่แท้งบุตร ซึ่งมีโอกาสเป็นไปได้ตั้งแต่ศูนย์จนถึงห้าครั้ง ข้อมูลนี้จึงควรจัดอยู่ในกลุ่มข้อมูลเชิงคุณภาพ และใช้หลักการวิเคราะห์ข้อมูลแบบข้อมูลเชิงคุณภาพ จึงจะเหมาะสมกว่า

2.2 Continuous numerical data หมายถึง ข้อมูลตัวเลขที่มีค่าได้ทุกค่าบนหน่วยของการวัด ซึ่งถือว่าเป็นข้อมูลที่มีความละเอียดมากที่สุด เช่น ระดับความดันโลหิต ระดับไขมันในเลือด ข้อมูลน้ำหนักและส่วนสูงของคนไข้ ซึ่งข้อมูลเชิงปริมาณเหล่านี้สามารถนำมาจัดเป็นข้อมูลกลุ่มได้ในภายหลัง ซึ่งจะเห็นได้ว่าข้อมูลที่มีความละเอียดมาก ๆ สามารถทำให้มีความละเอียดลดลงได้ แต่ข้อมูลที่มีความละเอียดน้อย เช่น ข้อมูลกลุ่ม จะทำให้มีความละเอียดเพิ่มขึ้นไม่ได้ ดังนั้นในขั้นตอนการเก็บข้อมูล ผู้วิจัยจึงควรพิจารณาว่าชนิดของข้อมูลนั้นๆเป็นอย่างไร ถ้าข้อมูลนั้นโดยธรรมชาติเป็นข้อมูลที่มีความละเอียดมากอยู่แล้ว เช่น อายุ เราควรที่จะเก็บข้อมูลนั้นตามลักษณะของข้อมูลที่วัดได้จริง แล้วจึงนำมาจัดกลุ่มในภายหลัง ในทางตรงกันข้าม ถ้าเราเก็บข้อมูลอายุ เป็นแบบกลุ่ม เช่น กลุ่มที่ 1 อายุน้อยกว่า 60 / กลุ่มที่ 2 อายุมากกว่า 60 ขึ้นไป ต่อมาภายหลังต้องการทราบรายละเอียดของข้อมูล หรือต้องการจัดกลุ่มใหม่ จะไม่สามารถทำได้ด้วยข้อมูลที่มีอยู่



รูปที่ 1 ชนิดของข้อมูลจำแนกตามลักษณะการวัด

บทที่ 2

การสรุปและนำเสนอข้อมูล (Summarizing data)

การสรุปข้อมูลเป็นขั้นตอนแรกในการวิเคราะห์ข้อมูล ซึ่งจะให้เห็นภาพรวมของข้อมูลที่ได้ว่ามีลักษณะอย่างไร ซึ่งสามารถทำได้โดยใช้สถิติเชิงพรรณนา (Descriptive statistics) ทั้งนี้การสรุปและนำเสนอข้อมูลขึ้นอยู่กับชนิดของข้อมูลนั้นๆว่าเป็น ข้อมูลเชิงคุณภาพ (ข้อมูลกลุ่ม) หรือข้อมูลเชิงปริมาณ (ข้อมูลต่อเนื่อง)

การสรุปและนำเสนอข้อมูลกลุ่ม (Categorical data)

ข้อมูลกลุ่ม คือข้อมูลที่มีการวัดแบบ nominal / ordinal scale เช่น เพศ กลุ่มเลือด สถานภาพสมรส เป็นต้น ข้อมูลชนิดนี้นำเสนอด้วย จำนวน และ ร้อยละ จำนวนในที่นี้หมายถึง จำนวนตัวอย่างที่ตกอยู่ในแต่ละกลุ่ม เมื่อถูกหารด้วยจำนวนตัวอย่างทั้งหมดที่มีอยู่ จะได้เป็นค่าร้อยละของจำนวนนั้นๆ ซึ่งผู้วิจัยควรนำเสนอทั้งสองค่าพร้อมกัน ไม่ควรนำเสนอค่าใดค่าหนึ่ง เพราะอาจทำให้การแปลผลผิดได้

ตัวอย่างที่ 1 ในการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิด Sleep apnea (SA) โดยจุลวิทย์และคณะ เพื่อนำไปใช้ในการสร้างคะแนนสำหรับทำนายการเกิด SA ซึ่งมีรูปแบบการศึกษาแบบภาคตัดขวาง (cross-sectional study) โดยในขั้นตอนแรกผู้วิจัยจะต้องนำเสนอภาพรวมของกลุ่มตัวอย่างว่ามีลักษณะอย่างไร ซึ่งการนำเสนอภาพรวมของกลุ่มตัวอย่างที่มีลักษณะเป็นข้อมูลเชิงคุณภาพ เช่น เพศ สามารถทำได้โดยใช้โปรแกรม STATA ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Tables --> One-way tables** โปรแกรมจะแสดงผลลัพธ์ที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. tab sex
```

sex	Freq.	Percent	Cum.
Female	266	31.78	31.78
Male	571	68.22	100.00
Total	837	100.00	

จากผลลัพธ์ที่ได้ แสดงให้เห็นว่ามีจำนวนตัวอย่างทั้งสิ้น 837 คน โดยแบ่งออกเป็นผู้หญิง 266 คน คิดเป็น 31.78% ของตัวอย่างทั้งหมด และเป็นผู้ชาย 571 คน คิดเป็น 68.22% ของตัวอย่างทั้งหมด

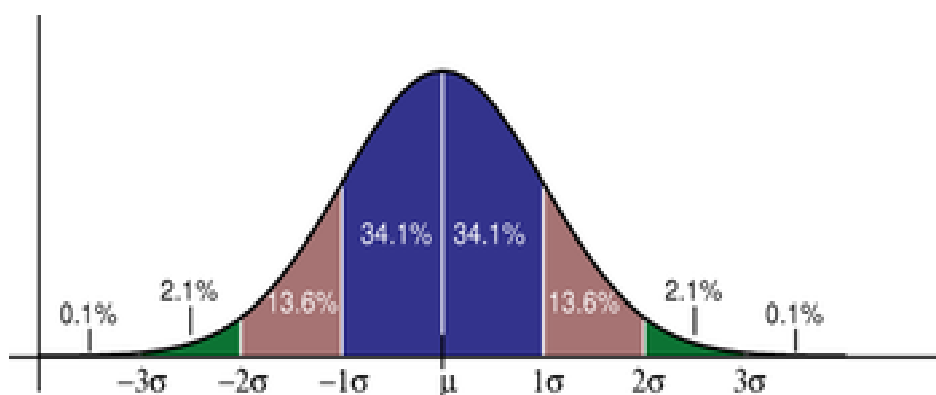
การสรุปและนำเสนอข้อมูลต่อเนื่อง (Continuous data)

ข้อมูลต่อเนื่อง คือข้อมูลที่มีการวัดแบบ discrete / continuous scale เช่น อายุ ความดันโลหิต น้ำหนัก ส่วนสูง เป็นต้น ข้อมูลชนิดนี้นำเสนอโดยใช้ค่ากลาง (measure of central tendency) เช่น ค่าเฉลี่ย (mean) ค่ามัธยฐาน (median) ร่วมกับค่าการกระจาย (measure of dispersion) เช่น ค่าส่วนเบี่ยงเบนมาตรฐาน (SD) ค่าต่ำสุดและสูงสุด (range) โดยมีหลักเกณฑ์ในการนำเสนอ ดังนี้

1. นำเสนอโดยใช้ค่าเฉลี่ย และค่าส่วนเบี่ยงเบนมาตรฐาน ในกรณีที่ข้อมูลมีการแจกแจงแบบปกติ
2. นำเสนอโดยใช้ค่ามัธยฐาน และค่าต่ำสุด-สูงสุด ในกรณีที่ข้อมูลไม่ใช้การแจกแจงแบบปกติ

ลักษณะของข้อมูลที่มีการแจกแจงแบบปกติ

1. ลักษณะการกระจายของข้อมูลเป็นแบบโค้งปกติ (รูประฆังคว่ำ) โดยที่ข้อมูลส่วนใหญ่กระจายตัวอยู่ตรงกลาง (95% ของข้อมูลกระจายอยู่ในช่วง $\text{mean} \pm 2\text{SD}$) และมีค่าสูงและต่ำกระจายตัวน้อยลงไปตามลำดับ (รูปที่ 2)
2. ค่าเฉลี่ย ค่ามัธยฐาน และค่าฐานนิยม อยู่ตำแหน่งเดียวกัน
3. สัมประสิทธิ์ความโด่งของการแจกแจง (Kurtosis) มีค่าเท่ากับ 3 และสัมประสิทธิ์ความเบ้ของการแจกแจง (Skewness) มีค่าเท่ากับ 0
4. หลักการคร่าวๆในการพิจารณาว่าข้อมูลมีการแจกแจงแบบปกติหรือไม่ ให้พิจารณาจากค่าส่วนเบี่ยงเบนมาตรฐาน โดยข้อมูลที่มีการแจกแจงแบบปกติ จะมีค่าส่วนเบี่ยงเบนมาตรฐานต่ำกว่าครึ่งหนึ่งของค่าเฉลี่ย แต่ถ้าส่วนเบี่ยงเบนมาตรฐานมีค่าสูงกว่าครึ่งหนึ่งของค่าเฉลี่ย มักจะพบว่าข้อมูลไม่ได้มีการแจกแจงแบบปกติ หรือข้อมูลมีการกระจายตัวแบบเบ้ (skew) ไปในทิศทางใดทิศทางหนึ่งนั่นเอง



รูปที่ 2 การแจกแจงแบบโค้งปกติ

การคำนวณหาค่ากลาง (Measures of central tendency)

1) ค่าเฉลี่ย

เป็นการวัดค่ากลางที่ใช้เป็นตัวแทนของตัวอย่างทั้งหมด ซึ่งสามารถหาได้จากผลรวมของค่าที่วัดได้จากตัวอย่างทั้งหมด หารด้วยขนาดของตัวอย่าง ดังสมการ (1)

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n} \quad (1)$$

เมื่อ $\bar{x}$ คือ ค่าเฉลี่ยของตัวอย่าง

n คือ ขนาดตัวอย่าง

x_i คือ ค่าที่วัดได้จากตัวอย่างแต่ละคน

$\sum$ คือ เครื่องหมายรวม (summation)

2) ค่ามัธยฐาน

เป็นอีกค่าหนึ่งที่ใช้เป็นค่ากลางของตัวอย่าง โดยเมื่อนำข้อมูลมาเรียงลำดับจากน้อยไปมาก หรือมากไปน้อย ค่ามัธยฐานคือ ค่าที่อยู่ตำแหน่งตรงกลางของข้อมูล ซึ่งสามารถหาได้จาก $(n+1)/2$ เมื่อ n คือจำนวนตัวอย่าง ยกตัวอย่างเช่น ถ้าจำนวนตัวอย่างเท่ากับ 11 ค่ามัธยฐาน คือค่าที่อยู่ตรงตำแหน่งที่ $(11+1)/2=6$ ของข้อมูล

ที่เรียงลำดับ ในทำนองเดียวกันถ้าจำนวนตัวอย่างเท่ากับ 12 ค่ามัธยฐาน คือค่าที่อยู่ตรงตำแหน่ง $(12+1)/2=6.5$ ซึ่งก็คือค่าเฉลี่ยของข้อมูลที่เรียงลำดับตรงตำแหน่งที่ 6 กับ 7

ตัวอย่างที่ 2 ข้อมูลอายุของผู้ป่วยโรคไต จำนวน 9 คน

ลำดับ	1	2	3	4	5	6	7	8	9
อายุ	23	28	28	31	32	34	37	42	50

ค่ามัธยฐาน คือค่าที่อยู่ตรงตำแหน่งที่ $(9+1)/2=5$ ของข้อมูลที่เรียงลำดับ ซึ่งก็คือ 32

ตัวอย่างที่ 3 ข้อมูลอายุของผู้ป่วยโรคไต จำนวน 10 คน

ลำดับ	1	2	3	4	5	6	7	8	9	10
อายุ	23	28	28	31	32	34	37	42	50	61

ค่ามัธยฐาน คือค่าที่อยู่ตรงตำแหน่งที่ $(10+1)/2=5.5$ ของข้อมูลที่เรียงลำดับ ซึ่งก็คือ $(32+34)/2=33$

การคำนวณค่าการกระจาย (Measures of dispersion)

1) ค่าส่วนเบี่ยงเบนมาตรฐาน

เป็นค่าที่แสดงความแปรปรวนของข้อมูลเมื่อเทียบกับค่าเฉลี่ย ถ้าส่วนเบี่ยงเบนมาตรฐานมีค่าสูง แสดงว่าข้อมูลที่วัดได้มีความแตกต่างกันมาก ในทางกลับกันถ้าส่วนเบี่ยงเบนมาตรฐานมีค่าต่ำ แสดงว่าข้อมูลส่วนใหญ่มีค่าใกล้เคียงกัน คำนี้อาจคำนวณได้จากสมการ (2)

$$SD = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1}} \quad (2)$$

เมื่อ $\bar{x}$ คือ ค่าเฉลี่ยของตัวอย่าง

n คือ ขนาดตัวอย่าง

x_i คือ ค่าที่วัดได้จากตัวอย่างแต่ละคน

$\sum$ คือ เครื่องหมายรวม (summation)

2) ค่าต่ำสุด และสูงสุด

ค่าต่ำสุด และค่าสูงสุดของข้อมูล หรือเรียกว่า ขอบเขตของข้อมูล เป็นค่าที่บ่งบอกถึงความแปรปรวนของข้อมูลอีกค่าหนึ่งที่น่าสนใจ

ตัวอย่างที่ 4 จากการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิด SA ในตัวอย่างที่ 1 ผู้วิจัยต้องนำเสนอภาพรวมของกลุ่มตัวอย่างที่มีลักษณะเป็นข้อมูลเชิงปริมาณ เช่น อายุ และ ESS score ซึ่งสามารถทำได้โดยใช้โปรแกรม STATA ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Summary statistics --> Summary statistics**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

1) ข้อมูลอายุ

```
. summarize age, detail
```

age					

	Percentiles	Smallest			
1%	19.8137	18.02466			
5%	25.08219	18.18356			
10%	30.55342	18.34521	Obs	837	
25%	38.83014	18.44384	Sum of Wgt.	837	
50%	49.92329	(2)	Mean	49.56676	(1)
		Largest	Std. Dev.	14.30717	
75%	60.23561	82.16164	Variance	204.6951	
90%	68.62466	82.18082	Skewness	-.0609356	
95%	72.27671	82.90137	Kurtosis	2.327515	
99%	78.85206	84.75616			

จากผลลัพธ์ที่ได้แสดงว่า ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐาน (1) จากตัวอย่างนี้มีค่าเท่ากับ 49.57 และ 14.31 ตามลำดับ ส่วนค่ามัธยฐานคือ ค่าที่อยู่ตรงตำแหน่ง Percentile ที่ 50 (2) ซึ่งจากตัวอย่างนี้ค่ามัธยฐานคือ 49.92 ค่าต่ำสุดคือ 18.02 และค่าสูงสุดคือ 84.76 เมื่อพิจารณาส่วนเบี่ยงเบนมาตรฐานของข้อมูลชุดนี้พบว่ามีค่าต่ำกว่าครึ่งหนึ่งของค่าเฉลี่ย และค่ามัธยฐานใกล้เคียงกับค่าเฉลี่ย โดยที่มีค่าความเบ้เท่ากับ -0.06 แสดงว่าข้อมูลมีการแจกแจงแบบปกติ เพราะฉะนั้นการสรุปและนำเสนอข้อมูลชุดนี้จึงใช้ค่าเฉลี่ย ร่วมกับส่วนเบี่ยงเบนมาตรฐาน

2) ข้อมูล ESS score

```
. summarize ess, detail
```

ess					

	Percentiles	Smallest			
1%	0	0			
5%	1	0			
10%	3	0	Obs	837	
25%	7	0	Sum of Wgt.	837	
50%	10		Mean	10.69415	
		Largest	Std. Dev.	5.747446	
75%	15	24	Variance	33.03314	
90%	19	24	Skewness	.1055136	
95%	20	24	Kurtosis	2.278389	
99%	23	24			

จากผลลัพธ์ที่ได้แสดงว่า ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐาน จากตัวอย่างนี้มีค่าเท่ากับ 10.69 และ 5.75 ตามลำดับ ส่วนค่ามัธยฐานคือ ค่าที่อยู่ตรงตำแหน่ง Percentile ที่ 50 ซึ่งจากตัวอย่างนี้ค่ามัธยฐานคือ 10 ค่าต่ำสุดคือ 0 และค่าสูงสุดคือ 24 และเมื่อพิจารณาส่วนเบี่ยงเบนมาตรฐานของข้อมูลชุดนี้พบว่ามีค่าเกินครึ่งหนึ่งของค่าเฉลี่ย และค่ามัธยฐานที่ต่ำกว่าค่าเฉลี่ย โดยที่มีค่าความเบ้เท่ากับ 0.11 แสดงว่าข้อมูลไม่ได้มีการแจกแจงแบบปกติ เพราะฉะนั้นการสรุปและนำเสนอข้อมูลชุดนี้จึงใช้ค่ามัธยฐาน ร่วมกับค่าต่ำสุดและสูงสุดของข้อมูล

บทที่ 3

การอ้างอิงทางสถิติ (Statistical inference)

ในการทำงานวิจัยส่วนใหญ่ เราทำการศึกษาจากกลุ่มตัวอย่าง (Sample) ที่ถูกเลือกมาอย่างสุ่มจากประชากร (Population) ทั้งนี้เพื่อเป็นการลดต้นทุนในด้านค่าใช้จ่าย เวลา และ กำลังคนที่ใช้ในการทำงานวิจัย โดยเราจะศึกษาลักษณะที่สนใจจากกลุ่มตัวอย่าง และนำค่าที่ได้จากกลุ่มตัวอย่างอ้างอิงไปยังประชากร ซึ่งค่าที่ได้จากกลุ่มตัวอย่างเรียกว่า ค่าสถิติ (Sample statistics) ส่วนลักษณะประชากรที่สนใจเรียกว่า ค่าพารามิเตอร์ (Parameter)

การอ้างอิงค่าสถิติไปยังประชากรแบ่งเป็น 2 รูปแบบตามวัตถุประสงค์หลักของงานวิจัย คือ การประมาณค่าพารามิเตอร์ และการทดสอบสมมติฐาน

การประมาณค่าพารามิเตอร์ (Parameter estimation)

เป็นการประมาณค่าลักษณะของประชากรที่สนใจจากกลุ่มตัวอย่างที่ถูกคัดเลือกมาเป็นตัวแทนของประชากร โดยการประมาณค่าแบ่งออกเป็น 2 ประเภทคือ การประมาณค่าแบบจุด (point estimation) และการประมาณค่าแบบช่วง (range estimation)

การประมาณค่าแบบจุด

เป็นการประมาณค่าลักษณะของประชากรที่สนใจจากกลุ่มตัวอย่างที่ถูกคัดเลือกให้เป็นตัวแทนของประชากร โดยการประมาณค่าแบ่งตามชนิดของข้อมูลดังนี้

1. การประมาณค่าสัดส่วน สำหรับข้อมูลกลุ่ม เช่น การประมาณค่าความชุกของภาวะความดันโลหิตสูงในประชากรไทย ซึ่งมีค่าเท่ากับ 27.6 % การประมาณค่าสัดส่วนการใช้จ่ายเงินโบราณในประชากรภาคอีสาน ซึ่งมีค่าเท่ากับ 33.3%
2. การประมาณค่าเฉลี่ย สำหรับข้อมูลต่อเนื่อง เช่น การประมาณค่าเฉลี่ยของระดับไขมันในเลือดในผู้ป่วยสูงอายุ ซึ่งมีค่าเท่ากับ 205.2 mg/dl การประมาณค่าเฉลี่ยอายุของผู้ป่วยที่เป็นโรคไตเรื้อรังในประเทศไทย ซึ่งมีค่าเท่ากับ 56.2 ปี

การประมาณค่าแบบช่วง

เป็นการประมาณค่าขอบเขตความเป็นไปได้ของลักษณะประชากรที่สนใจ โดยที่มีค่าจริงของลักษณะประชากรตกอยู่ในขอบเขตนั้นๆ การประมาณค่าแบบช่วงจะถูกกำหนดด้วยระดับความเชื่อมั่นเสมอ จึงเรียกค่าประมาณนี้ว่า ช่วงความเชื่อมั่น (Confidence interval)

วิธีการประมาณค่าแบบช่วงต้องอาศัยหลักเกณฑ์ของการแจกแจงแบบสุ่ม (Sampling distribution) ซึ่งการประมาณค่าจะแบ่งตามชนิดของข้อมูลดังนี้

1) การประมาณค่าช่วงความเชื่อมั่นของค่าสัดส่วน สำหรับข้อมูลกลุ่ม

เนื่องจากค่าสัดส่วนมีการแจกแจงแบบปกติ เพราะฉะนั้นการประมาณค่าช่วงความเชื่อมั่นของค่าสัดส่วนจึงถูกกำหนดด้วยค่า $z_{(1-\alpha/2)}$ ซึ่งค่าช่วงความเชื่อมั่นของค่าสัดส่วนสามารถคำนวณได้จากสูตรดังนี้

$$p \pm z_{(1-\alpha/2)} \times SE$$

เมื่อ

- p คือ ค่าจุดประมาณของสัดส่วน

- SE ของค่าสัดส่วนคือ $SE = \sqrt{p(1-p)/n}$
- $z_{(1-\alpha/2)}$ คือ ค่าที่ได้จากตาราง standard normal distribution ซึ่งค่านี้ขึ้นอยู่กับค่า α ดังนี้

α	$z_{(1-\alpha/2)}$
0.01	2.576
0.05	1.960
0.10	1.645

ตัวอย่างที่ 1 จากการสุ่มตัวอย่างของประชากรไทยจำนวน 3,459 คน และทำการซักประวัติ พร้อมทั้งตรวจร่างกาย พบว่า 955 คน มีประวัติความดันโลหิตสูง และ/หรือ systolic blood pressure มากกว่าหรือเท่ากับ 140 และ/หรือ diastolic blood pressure มากกว่าหรือเท่ากับ 90 ดังนั้นค่าประมาณของสัดส่วนผู้ป่วยที่มีภาวะความดันโลหิตสูง เท่ากับ $955/3,459=0.276$ และ 95% ช่วงความเชื่อมั่นของค่าสัดส่วนภาวะความดันโลหิตสูงในประชากรไทยมีค่า เท่ากับ

$$0.276-(1.96 \times 0.0076)=\mathbf{0.261} \text{ ถึง } 0.276+(1.96 \times 0.0076)=\mathbf{0.291}$$

โดยที่ค่า SE สามารถคำนวณได้จาก

$$SE = \sqrt{\frac{0.276(1-0.276)}{3459}} = 0.0076$$

นั่นหมายความว่า มีความเชื่อมั่น 95% ที่ค่าสัดส่วนของภาวะความดันโลหิตสูงในประชากรไทย มีค่าอยู่ระหว่าง 26% ถึง 29%

2) การประมาณค่าช่วงความเชื่อมั่นของค่าเฉลี่ย สำหรับข้อมูลต่อเนื่อง

เนื่องจากค่าเฉลี่ยมีการแจกแจงแบบ t distribution เพราะฉะนั้นการประมาณค่าช่วงความเชื่อมั่นของค่าเฉลี่ย จึงถูกกำหนดโดยค่า $t_{(1-\alpha/2)}$ ซึ่งค่าช่วงความเชื่อมั่นของค่าเฉลี่ยสามารถคำนวณได้จากสูตรดังนี้

$$\bar{x} \pm t_{(1-\alpha/2)} \times SE$$

เมื่อ

- $\bar{x}$ คือ ค่าจุดประมาณของค่าเฉลี่ย
- SE ของค่าเฉลี่ยคือ $SE = SD/\sqrt{n}$
- $t_{(1-\alpha/2)}$ คือค่าที่ได้จากตาราง t distribution ที่ degree of freedom เท่ากับ $n-1$ ซึ่งค่านี้จะแปรผันตามค่า α

และขนาดของตัวอย่าง เช่น

Sample size (n)	df	α	$t_{(1-\alpha/2)}$
50	49	0.01	2.680
50	49	0.05	2.010
50	49	0.10	1.677
100	99	0.01	2.626
100	99	0.05	1.984
100	99	0.10	1.660

ตัวอย่างที่ 2 จากการสุ่มตัวอย่างของประชากรสูงอายุจำนวน 167 คน และทำการตรวจร่างกาย พบว่า ค่าเฉลี่ยของระดับไขมันในเลือดเท่ากับ 206.47 mg/dl และมีค่าส่วนเบี่ยงเบนมาตรฐานเท่ากับ 42.19 โดยที่ SE ของค่าเฉลี่ยมีค่าเท่ากับ $42.19/\sqrt{167} = 3.26$ เมื่อกำหนด $\alpha = 0.05$ ค่า $t_{(1-\alpha/2)}$ ที่ระดับ degree of freedom 166 มีค่าเท่ากับ 1.974 (ตารางที่ 2 ใน Gardner MJ and Altman DG) ดังนั้น 95% ช่วงความเชื่อมั่นของค่าเฉลี่ยระดับไขมันในเลือดในประชากรสูงอายุมีค่าเท่ากับ

$$206.47 - (1.974 \times 3.26) = 200.03 \text{ ถึง } 206.47 + (1.974 \times 3.26) = 212.91$$

นั่นหมายความว่า มีความเชื่อมั่น 95% ที่ค่าเฉลี่ยของระดับไขมันในเลือดในประชากรสูงอายุ มีค่าอยู่ระหว่าง 200.03 ถึง 212.91 mg/dl

อย่างไรก็ตามการรายงานผลการศึกษาลักษณะของประชากรที่สนใจ ควรนำเสนอค่าจุดประมาณ (Point estimate) ร่วมกับช่วงความเชื่อมั่น (Confidence interval) เสมอ เพื่อแสดงให้เห็นความแม่นยำของจุดประมาณว่า น่าเชื่อถือมากน้อยเพียงใด ซึ่งความแม่นยำของจุดประมาณจะขึ้นอยู่กับขนาดของตัวอย่างเป็นสำคัญ โดยที่ขนาดตัวอย่างที่ใหญ่ จะให้ช่วงความเชื่อมั่นที่แคบ นั้นแสดงว่าค่าจุดประมาณมีความแม่นยำและน่าเชื่อถือ ในทางตรงกันข้าม ขนาดตัวอย่างที่เล็ก จะให้ช่วงความเชื่อมั่นที่กว้าง แสดงว่าจุดประมาณมีความแม่นยำต่ำและไม่น่าเชื่อถือ

การทดสอบสมมติฐาน (Hypothesis testing)

เป็นการทดสอบความแตกต่างของประชากร ซึ่งอาจเป็นการทดสอบสมมติฐานในประชากรกลุ่มเดียวหรือทดสอบในประชากรหลายกลุ่ม ยกตัวอย่างเช่น

- การทดสอบความแตกต่างค่าความชุกของโรคเอดส์ในปี 2005 กับเมื่อ 5 ปีที่ผ่านมา
- การทดสอบความแตกต่างของผลการรักษาของผู้ป่วย Myasthenia gravis ที่ได้รับการรักษาด้วยการผ่าตัดต่อมไทมัส (Thymectomy) กับผู้ป่วยที่ได้รับการรักษาด้วยยาเพียงอย่างเดียว (Non-thymectomy)
- การทดสอบความแตกต่างค่าเฉลี่ยไขมันในเลือดของผู้ป่วยที่ได้รับยา A กับผู้ป่วยที่ได้รับยา B

การตั้งสมมติฐานเพื่อการทดสอบแบ่งออกเป็น

1. สมมติฐานนัล (Null hypothesis) เป็นสมมติฐานที่ตั้งขึ้นในลักษณะที่ไม่มีความแตกต่างกัน เช่น
 - ค่าเฉลี่ยความหนาแน่นของกระดูกในกลุ่มที่ได้รับแคลเซียมเสริมไม่แตกต่างจากกลุ่มที่ได้รับยาหลอก ($H_0 : \mu_1 = \mu_2$)
 - อัตราการตายของผู้ป่วยมะเร็งตับที่ได้รับการรักษาด้วยการผ่าตัดไม่แตกต่างจากผู้ป่วยมะเร็งตับที่ได้รับการรักษาด้วยวิธีเคมีบำบัด ($H_0 : P_1 = P_2$)
2. สมมติฐานทางเลือก (Alternative hypothesis) เป็นสมมติฐานที่ตรงกันข้ามกับสมมติฐานนัล นั่นคือ มีความแตกต่างกัน เช่น
 - ค่าเฉลี่ยความหนาแน่นของกระดูกในกลุ่มที่ได้รับแคลเซียมเสริมแตกต่างจากกลุ่มที่ได้รับยาหลอก ($H_a : \mu_1 \neq \mu_2$)
 - อัตราการตายของผู้ป่วยมะเร็งตับที่ได้รับการรักษาด้วยการผ่าตัดแตกต่างจากผู้ป่วยมะเร็งตับที่ได้รับการรักษาด้วยวิธีเคมีบำบัด ($H_a : P_1 \neq P_2$)

ผลที่ได้จากการทดสอบสมมติฐาน คือ การปฏิเสธ หรือ ยอมรับสมมติฐาน ถ้าเราปฏิเสธสมมติฐานนัล แสดงว่าเรา ยอมรับสมมติฐานทางเลือก ในทางกลับกัน ถ้าเรายอมรับสมมติฐานนัล แสดงว่าเราปฏิเสธสมมติฐานทางเลือก

นอกจากนี้การตั้งสมมติฐานเพื่อการทดสอบยังสามารถแบ่งออกเป็น

1. การทดสอบสมมติฐานแบบทางเดียว (one-tailed test) เป็นการตั้งสมมติฐานแบบมีทิศทาง ซึ่งทิศทางอาจเป็นไปได้ทั้งในทางมากกว่า หรือน้อยกว่า การตั้งสมมติฐานลักษณะนี้ จะทำเมื่อสิ่งที่ต้องการทดสอบมีหลักฐานสนับสนุนที่แน่ชัดว่าความแตกต่างนั้นเป็นไปในทิศทางใด เช่น

- $H_0 : \mu_1 \geq \mu_2$ and $H_a : \mu_1 < \mu_2$
- $H_0 : \mu_1 \leq \mu_2$ and $H_a : \mu_1 > \mu_2$

2. การทดสอบสมมติฐานแบบสองทาง (two-tailed test) เป็นการตั้งสมมติฐานแบบไม่มีทิศทาง การตั้งสมมติฐานลักษณะนี้ จะทำเมื่อสิ่งที่ต้องการทดสอบไม่มีหลักฐานสนับสนุนที่ชัดเจนเกี่ยวกับทิศทางของความแตกต่าง เช่น

- $H_0 : \mu_1 = \mu_2$ and $H_a : \mu_1 \neq \mu_2$

ในทางปฏิบัติแนะนำให้ตั้งสมมติฐานแบบสองทางจะปลอดภัยที่สุดสำหรับนักวิจัย และทำให้ได้ขนาดตัวอย่างที่มากกว่า การตั้งสมมติฐานแบบทางเดียวอีกด้วย

ความผิดพลาดที่เกิดจากการทดสอบสมมติฐาน (Error)

ผลลัพธ์ที่ได้จากการทดสอบสมมติฐานในกลุ่มตัวอย่างอาจผิดหรือถูกก็ได้ เมื่อนำไปอ้างอิงกับประชากรจริง การยอมรับหรือปฏิเสธสมมติฐานนั้นย่อมมีความผิดพลาด(error) เกิดขึ้นเสมอ โดยความผิดพลาดที่เกิดขึ้นนี้สามารถแบ่งได้ 2 กรณี คือ

1. ความผิดพลาดชนิดแรก (type one error, false positive) คือความผิดพลาดที่เกิดขึ้นในกรณีที่สมมติฐานนั้นถูก แต่เราปฏิเสธสมมติฐานนี้ ตัวอย่างเช่น ต้องการทดสอบความแตกต่างของอายุเฉลี่ยระหว่างประชากร 2 กลุ่ม ซึ่งในประชากรจริงพบว่าไม่มีความแตกต่างของอายุเฉลี่ยระหว่างประชากรทั้งสองกลุ่ม แต่ผลการศึกษากับกลุ่มตัวอย่างกลับพบว่ามีความแตกต่างกัน ดังนั้นเรามีโอกาสผิดพลาดได้ α % และมีค่าความเชื่อมั่นเท่ากับ $1 - \alpha$ %
2. ความผิดพลาดชนิดที่สอง (type two error, false negative) คือความผิดพลาดที่เกิดขึ้นในกรณีที่สมมติฐานนั้นผิด แต่เรายอมรับสมมติฐานนี้ ตัวอย่างเช่น ในประชากรจริงพบว่ามีความแตกต่างของอายุเฉลี่ยระหว่างประชากรสองกลุ่ม แต่การศึกษากับกลุ่มตัวอย่างพบว่าไม่มีความแตกต่าง นั่นคือ เรามีโอกาสผิดพลาด β % และมีค่าความเชื่อถือ (power of test) เท่ากับ $1 - \beta$ %

ตารางที่ 1 ความผิดพลาดที่เกิดจากการทดสอบสมมติฐาน

Statistical decision based on sample	In population	
	H_0 True	H_0 False
Significance	α (False positive)	$1 - \beta$ (Power of test)
Non-significance	$1 - \alpha$ (Confidence)	β (False negative)

ขั้นตอนการทดสอบสมมติฐาน (Hypothesis testing steps)

1. ตั้งสมมติฐานนัย และสมมติฐานทางเลือก ให้ชัดเจนและสอดคล้องกับวัตถุประสงค์ของงานวิจัย ตัวอย่างเช่น ต้องการทดสอบความแตกต่างค่าเฉลี่ยระดับไขมันในเลือดระหว่างผู้ป่วยที่ได้รับการรักษา

ด้วยยา A กับผู้ป่วยที่ได้รับการรักษาด้วยยา B การตั้งสมมติฐานสามารถทำได้ดังนี้ สมมติฐานนัย คือ ค่าเฉลี่ยไขมันในเลือดระหว่างผู้ป่วยที่ได้รับยา A และยา B ไม่แตกต่างกัน

$$H_0 : \mu_A = \mu_B$$

สมมติฐานทางเลือก คือค่าเฉลี่ยไขมันในเลือดระหว่างผู้ป่วยที่ได้รับยา A และยา B แตกต่างกัน

$$H_a : \mu_A \neq \mu_B$$

- กำหนดค่าระดับนัยสำคัญของการทดสอบสมมติฐาน (significant level of the hypothesis test) ซึ่งก็คือค่าความผิดพลาดชนิดแรกนั่นเอง โดยทั่วไปจะกำหนดค่านี้เท่ากับ 0.05 ซึ่งหมายความว่ายอมให้เกิดความผิดพลาดจากการปฏิเสธสมมติฐานที่ถูกต้อง ได้เพียง 5% และมีความเชื่อมั่นในการยอมรับสมมติฐานที่ถูกต้อง 95% ถ้าต้องการความเชื่อมั่นในการยอมรับสมมติฐานที่ถูกต้อง 99% จะต้องกำหนดระดับนัยสำคัญของการทดสอบสมมติฐานให้มีค่าเท่ากับ 0.01
- เลือกสถิติที่เหมาะสมในการทดสอบสมมติฐาน โดยพิจารณาจากชนิดของตัวแปรและวัตถุประสงค์ของงานวิจัยเป็นหลัก ซึ่งจะได้กล่าวถึงรายละเอียดในหัวข้อถัดไป
- คำนวณค่าสถิติ และแปลงเป็นค่าความน่าจะเป็น (probability) หรือที่เรียกว่าค่า p-value ซึ่งหมายถึงโอกาสที่จะพบเหตุการณ์ที่สนใจทดสอบในประชากร เมื่อสมมติฐานนัยถูกต้อง เช่น ผลการทดสอบความแตกต่างของค่าเฉลี่ยระดับความดันโลหิตระหว่างกลุ่มตัวอย่างที่ได้รับยาใหม่กับกลุ่มตัวอย่างที่ได้รับยาเก่า ได้ค่า p-value เท่ากับ 0.34 หมายความว่า ถ้าสมมติฐานนัยในประชากรถูกจริง ในตัวอย่างที่ศึกษามีโอกาสที่ค่าเฉลี่ยความดันโลหิตในกลุ่มตัวอย่างทั้งสองกลุ่มจะไม่แตกต่างกัน 34% ดังนั้นถ้า p-value มีค่ามาก แสดงว่าโอกาสที่จะพบว่าสมมติฐานนัยถูกมีมาก เราจึงไม่ปฏิเสธสมมติฐานนัย ในทางตรงข้ามถ้า p-value มีค่าน้อย แสดงว่าโอกาสที่จะพบว่าสมมติฐานนัยถูกมีน้อย เราจึงปฏิเสธสมมติฐานนัย ซึ่งในปัจจุบันโปรแกรมสำเร็จรูปสำหรับการวิเคราะห์ข้อมูลส่วนใหญ่มักจะรายงานผลการทดสอบสมมติฐานทั้งในรูปของค่าสถิติ และค่า p-value
- สรุปผลการทดสอบสมมติฐาน โดยการนำค่า p-value ที่ได้ในขั้นตอนที่ 4 ไปเปรียบเทียบกับค่าระดับนัยสำคัญของการทดสอบสมมติฐาน (α) ถ้า p-value มีค่าน้อยกว่า หรือเท่ากับค่า α เราจะปฏิเสธสมมติฐานนัย ในทางตรงกันข้ามถ้า p-value มีค่ามากกว่าค่า α เราจะไม่ปฏิเสธสมมติฐานนัย

บทที่ 4

การทดสอบสมมติฐานสำหรับข้อมูลต่อเนื่อง (Hypothesis testing for continuous data)

การเลือกใช้สถิติสำหรับการทดสอบสมมติฐานมีหลักการดังนี้

1. พิจารณาวัตถุประสงค์หลักของงานวิจัยว่าต้องการศึกษาความแตกต่างในประชากรหนึ่งกลุ่ม สองกลุ่ม หรือมากกว่าสองกลุ่มขึ้นไป
2. พิจารณาลักษณะของตัวแปรว่าเป็นข้อมูลกลุ่ม หรือข้อมูลต่อเนื่อง
3. ในกรณีที่ลักษณะตัวแปรที่สนใจเป็นข้อมูลต่อเนื่อง ต้องพิจารณาลักษณะการกระจายของข้อมูลว่ามีการกระจายตัวแบบโค้งปกติ หรือมีการกระจายตัวแบบเบ้

ในบทนี้จะกล่าวถึงหลักการในการเลือกใช้สถิติสำหรับการทดสอบสมมติฐานเฉพาะในกรณีที่ต้องการทดสอบความแตกต่างของประชากรหนึ่งกลุ่ม และสองกลุ่ม โดยลักษณะตัวแปรที่สนใจเป็นข้อมูลต่อเนื่องเท่านั้น ส่วนการทดสอบสมมติฐานสำหรับข้อมูลกลุ่มจะกล่าวถึงในบทถัดไป

การทดสอบสมมติฐานในประชากรสองกลุ่มที่ไม่อิสระต่อกัน

เช่น การศึกษาผลการให้ยา hydrochlorothiazide ต่อความดันโลหิต โดยผู้วิจัยได้ทำการวัดความดันโลหิตก่อนและหลังได้รับยา ซึ่งวัตถุประสงค์ของการศึกษานี้ คือการทดสอบความแตกต่างของความดันโลหิตระหว่างก่อนและหลังได้รับยาลดความดัน ซึ่งการทดสอบสมมติฐานในประชากรสองกลุ่มที่ไม่อิสระต่อกันแบ่งเป็น 2 กรณีคือ

1. ลักษณะการกระจายของข้อมูลเป็นแบบโค้งปกติ (Normal distribution)

ในกรณีที่ข้อมูลทั้งสองกลุ่มมีลักษณะการกระจายเป็นแบบโค้งปกติ เราจะทำการทดสอบความแตกต่างของค่าเฉลี่ยระหว่างประชากรสองกลุ่มที่ไม่เป็นอิสระต่อกัน โดยสถิติที่ใช้ในการทดสอบคือ **paired t-test** ซึ่งมีข้อตกลงเบื้องต้นที่สำคัญคือ ข้อมูลของทั้งสองกลุ่มต้องมีการกระจายตัวเป็นแบบโค้งปกติ ดังนั้นก่อนที่จะตัดสินใจใช้ paired t-test ผู้วิจัยต้องทำการทดสอบการกระจายของข้อมูลทั้งสองกลุ่มก่อนว่ามีการกระจายแบบโค้งปกติหรือไม่ โดยใช้สถิติ เช่น Skewness and Kurtosis test หรือ Shapiro-wilk test หรือ Shapiro-Francia test ในการทดสอบการกระจายของข้อมูล หรือพิจารณาลักษณะการกระจายของข้อมูลจากค่าส่วนเบี่ยงเบนมาตรฐานดังที่กล่าวไว้แล้วในหัวข้อ การแจกแจงข้อมูลแบบปกติ ในบทที่ 2

ตัวอย่างที่ 1 วีระวัฒน์และคณะได้ทำการศึกษาเรื่องประสิทธิภาพของยา Highly active antiretroviral therapy (HAART) regimens ในผู้ป่วยที่ติดเชื้อ HIV โดยผู้วิจัยต้องการทดสอบความแตกต่างของปัจจัยต่างๆ เช่น น้ำหนัก ค่า CD4 ค่า Viral load ว่ามีการเปลี่ยนแปลงอย่างไรหลังจากได้รับยาต้านไวรัสกลุ่มนี้ ในตัวอย่างนี้จะทำการทดสอบน้ำหนักระหว่างก่อนและหลังได้รับยาต้านไวรัส โดยมีขั้นตอนในการทดสอบสมมติฐานดังนี้

ขั้นตอนแรก คือการทดสอบการกระจายของข้อมูลของทั้งสองกลุ่ม (ก่อนและหลังได้รับยา) ว่ามีการกระจายแบบโค้งปกติหรือไม่ โดยมีขั้นตอนในการทดสอบดังนี้

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือก
สมมติฐานนัย คือข้อมูลมีการกระจายแบบโค้งปกติ
สมมติฐานทางเลือก คือข้อมูลไม่ได้มีการกระจายแบบโค้งปกติ
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. ใช้สถิติ Shapiro-wilk test ในการทดสอบการกระจายของข้อมูลแต่ละกลุ่ม การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Distributional plots and tests --> Shapiro-Wilk normality test**

โปรแกรมจะแสดงผลพีธีที่ได้นบนหน้าต่าง Stata Results ดังนี้

```
. swilk bw0 bw12
```

Shapiro-Wilk W test for normal data					
Variable	Obs	W	V	z	Prob>z
bw0	121	0.98043	1.896	1.434	0.07580
bw12	121	0.98356	1.593	1.044	0.14817

- ผลการทดสอบการกระจายของข้อมูลน้ำหนักก่อนและหลังได้รับยาต้านไวรัส พบว่า p-value มีค่ามากกว่า 0.05 ทั้งก่อนและหลังได้รับยา แสดงว่า การกระจายของข้อมูลน้ำหนักทั้งก่อนและหลังได้รับยาต้านไวรัสมีการแจกแจงแบบโค้งปกติ เพราะฉะนั้นขั้นตอนต่อไปคือการทดสอบความแตกต่างของค่าเฉลี่ยน้ำหนักระหว่างก่อนและหลังได้รับยาโดยใช้สถิติ **paired t-test**

ขั้นตอนการทดสอบความแตกต่างของค่าเฉลี่ยน้ำหนักก่อนและหลังได้รับยาต้านไวรัส

- ตั้งสมมติฐานน้ลและสมมติฐานทางเลือกดังนี้
สมมติฐานน้ล คือค่าเฉลี่ยน้ำหนักระหว่างก่อนและหลังได้รับยาต้านไวรัสไม่แตกต่างกัน ($H_0 : \mu_{bef} = \mu_{aft}$)
สมมติฐานทางเลือก คือค่าเฉลี่ยน้ำหนักระหว่างก่อนและหลังได้รับยาต้านไวรัสแตกต่างกัน ($H_a : \mu_{bef} \neq \mu_{aft}$)
- กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
- ใช้สถิติ paired t-test ในการทดสอบความแตกต่างของค่าเฉลี่ยน้ำหนักระหว่างก่อนและหลังได้รับยาต้านไวรัส
การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Classical tests of hypotheses --> Mean-comparison test, paired data**

```
. ttest bw0 == bw12
```

Paired t test

Variable	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
bw0	121	54.30248	.8660592	9.526651	52.58774	56.01722
bw12	121	57.31322	.9380435	10.31848	55.45596	59.17048
diff	121	-3.010744	.4654596	5.120055	-3.932321	-2.089166

mean(diff) = mean(bw0 - bw12) t = -6.4683
Ho: mean(diff) = 0 degrees of freedom = 120

Ha: mean(diff) < 0 Ha: mean(diff) != 0 Ha: mean(diff) > 0
Pr(T < t) = 0.0000 Pr(|T| > |t|) = 0.0000 Pr(T > t) = 1.0000

- ผลการทดสอบความแตกต่างค่าเฉลี่ยน้ำหนักระหว่างก่อนและหลังได้รับยาต้านไวรัส พบว่า p-value มีค่า <0.001 แสดงว่า ค่าเฉลี่ยน้ำหนักระหว่างก่อนและหลังได้รับยาต้านไวรัสแตกต่างกันอย่างมีนัยสำคัญทางสถิติ โดยค่าเฉลี่ยและค่าส่วนเบี่ยงเบนมาตรฐานของน้ำหนักก่อนได้รับยาต้านไวรัสเท่ากับ 54.30 kg และ 9.53 kg

ตามลำดับ ส่วนค่าเฉลี่ยและค่าส่วนเบี่ยงเบนมาตรฐานของน้ำหนักหลังได้รับยาต้านไวรัสเท่ากับ 57.31 kg และ 10.32 kg ตามลำดับ สรุปว่าการได้รับยาต้านไวรัสมีผลทำให้น้ำหนักตัวของผู้ป่วยเพิ่มขึ้นอย่างมีนัยสำคัญทางสถิติ

2. ลักษณะการกระจายของข้อมูลไม่ใช่โค้งปกติ (Non-normal distribution)

ในกรณีที่ข้อมูลทั้งสองกลุ่มไม่ได้มีการแจกแจงแบบโค้งปกติ เราจะทำการทดสอบความแตกต่างของค่ามัธยฐานแทนค่าเฉลี่ย โดยสถิติที่ใช้ในการทดสอบความแตกต่างระหว่าง 2 กลุ่มที่ไม่เป็นอิสระต่อกัน และลักษณะการกระจายของข้อมูลไม่ได้เป็นแบบโค้งปกติคือ **Wilcoxon signed-ranks test** ซึ่งเป็นสถิติที่ไม่ได้คำนึงถึงการแจกแจงของข้อมูล เราเรียกการวิเคราะห์ข้อมูลแบบนี้ว่า Non-parametric ซึ่งเป็นวิธีที่เหมาะสมสำหรับข้อมูลที่มีความเบ้มาๆ

ตัวอย่างที่ 2 จากการศึกษาของวีระวัฒน์และคณะ ในตัวอย่างที่ 1 โดยในตัวอย่างนี้จะทำการทดสอบความแตกต่างของจำนวน CD4 ระหว่างก่อนและหลังได้รับยาต้านไวรัส โดยมีขั้นตอนในการทดสอบสมมติฐานดังนี้

ขั้นตอนแรก คือการทดสอบการกระจายของข้อมูลของทั้งสองกลุ่ม (ก่อนและหลังได้รับยา) ว่ามีการกระจายแบบโค้งปกติหรือไม่ โดยมีขั้นตอนในการทดสอบดังนี้

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือก
สมมติฐานนัย คือข้อมูลมีการกระจายแบบโค้งปกติ
สมมติฐานทางเลือก คือข้อมูลไม่ได้มีการกระจายแบบโค้งปกติ
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. ใช้สถิติ Shapiro-wilk test ในการทดสอบการกระจายของข้อมูลแต่ละกลุ่ม การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Distributional plots and tests --> Shapiro-Wilk normality test**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. swilk cd40 cd412
```

Shapiro-Wilk W test for normal data					
Variable	Obs	W	V	z	Prob>z
cd40	121	0.77137	22.156	6.944	0.00000
cd412	119	0.87373	12.065	5.577	0.00000

4. ผลการทดสอบการกระจายของจำนวน CD4 ก่อนและหลังได้รับยาต้านไวรัส พบว่า p-value มีค่าน้อยกว่า 0.05 ทั้งก่อนและหลังได้รับยา แสดงว่า การกระจายของจำนวน CD4 ทั้งก่อนและหลังได้รับยาต้านไวรัสไม่ได้มีการแจกแจงแบบโค้งปกติ เพราะฉะนั้นขั้นตอนต่อไปคือการทดสอบความแตกต่างของค่ามัธยฐานของจำนวน CD4 ระหว่างก่อนและหลังได้รับยาโดยใช้สถิติ **Wilcoxon signed-ranks test**

ขั้นตอนการทดสอบความแตกต่างของค่ามัธยฐาน CD4 ก่อนและหลังได้รับยาต้านไวรัส

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือกดังนี้
สมมติฐานนัยคือ ค่ามัธยฐาน CD4 ระหว่างก่อนและหลังได้รับยาไม่แตกต่างกัน ($H_0 : M_{bef} = M_{aft}$)
สมมติฐานทางเลือกคือ ค่ามัธยฐาน CD4 ระหว่างก่อนและหลังได้รับยาแตกต่างกัน ($H_a : M_{bef} \neq M_{aft}$)
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. สถิติที่ใช้ในการทดสอบคือ Wilcoxon signed-rank test การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Nonparametric tests of hypotheses -->**

Wilcoxon matched-pairs signed-rank test

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. signrank cd40 = cd412

Wilcoxon signed-rank test

      sign |      obs   sum ranks   expected
-----+-----
      positive |      7      256.5     3570
      negative |     112     6883.5     3570
      zero |      0         0         0
-----+-----
      all |     119      7140     7140

unadjusted variance 142205.00
adjustment for ties      -5.38
adjustment for zeros      0.00
-----
adjusted variance 142199.63

Ho: cd40 = cd412
      z =      -8.787
Prob > |z| =      0.0000
```

4. ผลการทดสอบแตกต่างของจำนวน CD4 ก่อนและหลังได้รับยาต้านไวรัส พบว่า p-value มีค่าน้อยกว่า 0.05 แสดงว่าการเปลี่ยนแปลงของจำนวน CD4 ระหว่างก่อนและหลังได้รับยาต้านไวรัส แตกต่างกันอย่างมีนัยสำคัญทางสถิติ โดยค่ามัธยฐานของจำนวน CD4 ก่อนได้รับยาต้านไวรัสมีค่าเท่ากับ 32 และมีค่าต่ำสุด-สูงสุดคือ 1-358 ส่วนค่ามัธยฐานของจำนวน CD4 หลังได้รับยาต้านไวรัสมีค่าเท่ากับ 132 และมีค่าต่ำสุด-สูงสุดคือ 3-726 ซึ่งการคำนวณหาค่ามัธยฐานและค่าต่ำสุด-สูงสุดโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

```
. sum cd40 cd412 if cd412!=.,detail

-----+-----
0 cd4
Percentiles      Smallest
1%              1          1
5%              2          1
10%             4          1   Obs          119
25%            14          1   Sum of wgt.  119

50%            32          1   Mean          64.09244
75%            80          1   Std. Dev.     74.6893
90%           180         284   Variance     5578.491
95%           236         314   Skewness      1.707546
99%           314         358   Kurtosis      5.518846
```

12 cd4				
Percentiles		Smallest		
1%	10	3		
5%	27	10		
10%	34	11	Obs	119
25%	69	15	Sum of Wgt.	119
50%	132		Mean	162.6975
		Largest	Std. Dev.	128.598
75%	219	488		
90%	332	493	Variance	16537.43
95%	449	560	Skewness	1.518691
99%	560	726	Kurtosis	5.86181

การทดสอบสมมติฐานในประชากรสองกลุ่มที่อิสระต่อกัน

เช่น การศึกษาผลของการให้แคลเซียมเสริมในสตรีวัยหมดประจำเดือน ขนาด 400 และ 800 มิลลิกรัมต่อปริมาณความหนาแน่นของกระดูก (Bone mineral density) โดยผู้วิจัยได้แบ่งผู้ป่วยออกเป็นสองกลุ่มโดยการสุ่มว่าผู้ป่วยรายใดจะได้แคลเซียมขนาด 400 หรือ 800 มิลลิกรัม ซึ่งวัตถุประสงค์ของการศึกษานี้คือ การทดสอบความแตกต่างของความหนาแน่นกระดูกระหว่างผู้ป่วยที่ได้รับแคลเซียม 400 และ 800 มิลลิกรัม ซึ่งการทดสอบสมมติฐานในประชากรสองกลุ่มที่อิสระต่อกันแบ่งเป็น 2 กรณีคือ

1. ลักษณะการกระจายของข้อมูลเป็นแบบโค้งปกติ (Normal distribution)

ในกรณีที่ข้อมูลทั้งสองกลุ่มมีลักษณะการกระจายเป็นแบบโค้งปกติ เราจะทำการทดสอบความแตกต่างของค่าเฉลี่ยระหว่างประชากรสองกลุ่มที่อิสระต่อกัน โดยสถิติที่ใช้ในการทดสอบคือ **student's t-test** ซึ่งมีข้อตกลงเบื้องต้นที่สำคัญคือ ข้อมูลของทั้งสองกลุ่มต้องมีการกระจายตัวเป็นแบบโค้งปกติ ดังนั้นก่อนที่จะตัดสินใจใช้ student's t-test ผู้วิจัยต้องทำการทดสอบการกระจายของข้อมูลทั้งสองกลุ่มก่อนว่ามีการกระจายแบบโค้งปกติหรือไม่ โดยใช้สถิติ เช่น Skewness and Kurtosis test หรือ Shapiro-wilk test หรือ Shapiro-Francia test ในการทดสอบการกระจายของข้อมูล หรือพิจารณาลักษณะการกระจายของข้อมูลจากค่าส่วนเบี่ยงเบนมาตรฐานดังที่กล่าวไว้แล้วในหัวข้อ การแจกแจงข้อมูลแบบปกติ ในบทที่ 2

การทดสอบความแตกต่างของค่าเฉลี่ยโดยใช้ Student's t test แบ่งออกเป็น 2 กรณีคือ

กรณีที่ 1 เมื่อความแปรปรวนระหว่างประชากร 2 กลุ่มเท่ากัน จะทำการทดสอบความแตกต่างของค่าเฉลี่ยโดยใช้ Student's t-test แบบ pooled variance

กรณีที่ 2 เมื่อความแปรปรวนระหว่างประชากร 2 กลุ่มไม่เท่ากัน จะทำการทดสอบความแตกต่างของค่าเฉลี่ยโดยใช้ Student's t-test แบบ separated variance

ดังนั้นก่อนการทดสอบความแตกต่างของค่าเฉลี่ยโดย Student's t test ผู้วิจัยต้องทำการทดสอบความแปรปรวนระหว่างตัวอย่างทั้ง 2 กลุ่มก่อนว่าความแตกต่างกันหรือไม่ โดยสถิติที่ใช้ในการทดสอบคือ F-test

ตัวอย่างที่ 3 ในการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิด Sleep apnea (SA) โดยจุลวิทย์และคณะ เพื่อนำไปใช้ในการสร้างคะแนนสำหรับทำนายการเกิด SA ซึ่งมีรูปแบบการศึกษาแบบภาคตัดขวาง (cross-sectional study) ดังนั้นก่อนที่จะทำการวิเคราะห์เพื่อหาปัจจัยที่มีความสัมพันธ์กับการเกิด SA ผู้วิจัยจะต้องทดสอบความแตกต่างของปัจจัยต่าง ๆ กับการเกิด SA โดยพิจารณาที่ละปัจจัย ถ้าปัจจัยใดมีความแตกต่างกันระหว่างกลุ่มคนไข้ที่เป็น SA กับคนไข้ที่ไม่ได้เป็น SA แสดงว่าปัจจัยนั้นมีความสัมพันธ์กับการเกิด SA เราเรียกการทดสอบลักษณะนี้ว่า *Univariate analysis* ในที่ตัวอย่างจะแสดงการทดสอบความแตกต่างค่าเฉลี่ยอายุระหว่างกลุ่มผู้ป่วยที่เป็น SA กับกลุ่มผู้ป่วยที่ไม่เป็น SA โดยมีขั้นตอนในการทดสอบสมมติฐานดังนี้

ขั้นตอนแรก คือการทดสอบการกระจายของข้อมูลของทั้งสองกลุ่ม (กลุ่มผู้ป่วยที่เป็น SA และกลุ่มผู้ป่วยที่ไม่ได้เป็น SA) ว่ามีการกระจายแบบโค้งปกติหรือไม่ โดยมีขั้นตอนในการทดสอบดังนี้

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือก
สมมติฐานนัย คือข้อมูลมีการกระจายแบบโค้งปกติ
สมมติฐานทางเลือก คือข้อมูลไม่ได้มีการกระจายแบบโค้งปกติ
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. ใช้สถิติ Shapiro-wilk test ในการทดสอบการกระจายของข้อมูลแต่ละกลุ่ม การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Distributional plots and tests --> Shapiro-Wilk normality test**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. by SA, sort : swilk age
```

```
-> SA = 0
```

Shapiro-Wilk W test for normal data

Variable	Obs	W	V	z	Prob>z
age	268	0.97765	4.309	3.410	0.00032

```
-> SA = 1
```

Shapiro-Wilk W test for normal data

Variable	Obs	W	V	z	Prob>z
age	569	0.99270	2.759	2.454	0.00707

4. ผลการทดสอบการกระจายของข้อมูลอายุในกลุ่มผู้ป่วยที่เป็น SA และกลุ่มผู้ป่วยที่ไม่ได้เป็น SA พบว่า p-value มีค่าน้อยกว่า 0.05 ทั้งสองกลุ่ม แสดงว่า การกระจายของข้อมูลอายุของทั้งสองกลุ่มไม่ได้มีการแจกแจงแบบโค้งปกติ แต่เนื่องจากสถิติที่ใช้ในการทดสอบการกระจายของข้อมูลมีความไวมากเมื่อตัวอย่างมีขนาดใหญ่ นั้นหมายความว่าข้อมูลที่มีความเบ้เล็กน้อย ตัวสถิติก็สามารถตรวจพบได้ เพราะฉะนั้นในกรณีที่ตัวอย่างมีขนาดค่อนข้างใหญ่ควรพิจารณาการกระจายของข้อมูลว่ามีการแจกแจงแบบโค้งปกติหรือไม่ด้วยค่าสัมประสิทธิ์ต่างๆ เช่น ค่าความโด่ง (kurtosis) ค่าความเบ้ (skewness) หรือค่าส่วนเบี่ยงเบนมาตรฐาน โดยใช้โปรแกรม STATA ในการคำนวณค่าต่างๆดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Summary and descriptive statistics --> Summary statistics**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. by SA, sort: sum age,detail
```

-> SA = 0

age				
	Percentiles	Smallest		
1%	18.44384	18.02466		
5%	20.89863	18.18356		
10%	24.27945	18.44384	Obs	268
25%	31.97123	18.83014	Sum of Wgt.	268
50%	41.99178		Mean	43.15004
		Largest	Std. Dev.	14.61313
75%	53.52877	75.98356		
90%	64.03835	75.9863	Variance	213.5436
95%	68.38356	80.03561	Skewness	.3083409
99%	75.9863	81.19452	Kurtosis	2.386439

-> SA = 1

age				
	Percentiles	Smallest		
1%	23.84384	18.34521		
5%	30.63836	19.47945		
10%	34.87945	20.02192	Obs	569
25%	42.67945	22.38356	Sum of Wgt.	569
50%	53.10137		Mean	52.58904
		Largest	Std. Dev.	13.12702
75%	62.20822	82.16164		
90%	70.27123	82.18082	Variance	172.3185
95%	73.07397	82.90137	Skewness	-.1092881
99%	80.4	84.75616	Kurtosis	2.423948

5. จากการพิจารณาลักษณะการกระจายของข้อมูลทั้งสองกลุ่มพบว่าค่าส่วนเบี่ยงเบนมาตรฐานไม่มากกว่าครึ่งหนึ่งของค่าเฉลี่ย และค่าสัมประสิทธิ์ความเบ้เข้าใกล้ศูนย์ ซึ่งเป็นไปตามลักษณะการแจกแจงแบบโค้งปกติ ส่วนค่าสัมประสิทธิ์ความโด่งมีค่าต่ำกว่า 3 ซึ่งไม่เป็นไปตามลักษณะการแจกแจงแบบโค้งปกติ อย่างไรก็ตามในขั้นตอนนี้เป็นทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยโดยเฉพาะ ซึ่งค่าสัมประสิทธิ์ที่มีผลกระทบต่อค่าเฉลี่ยคือค่าสัมประสิทธิ์ความเบ้ เพราะฉะนั้นในกรณีนี้เราไม่จำเป็นต้องพิจารณาค่าสัมประสิทธิ์ความโด่งของข้อมูล จากรายละเอียดดังกล่าวข้างต้น สรุปว่าข้อมูลอายุของทั้งสองกลุ่มมีการแจกแจงแบบโค้งปกติ และสามารถใช้สถิติ Student's t test ในการทดสอบความแตกต่างของค่าเฉลี่ยอายุระหว่างกลุ่มตัวอย่างสองกลุ่มในขั้นตอนต่อไปได้

ขั้นตอนการทดสอบความแตกต่างของค่าเฉลี่ยอายุระหว่างกลุ่มผู้ป่วย 2 กลุ่ม

ขั้นตอนที่ 1 ทดสอบความแตกต่างของความแปรปรวนระหว่างตัวอย่าง 2 กลุ่ม

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือกดังนี้

สมมติฐานนัยคือ ความแปรปรวนระหว่างประชากร 2 กลุ่ม ไม่แตกต่างกัน ($H_0 : \sigma_A^2 = \sigma_B^2$)

สมมติฐานทางเลือกคือ ความแปรปรวนระหว่างประชากร 2 กลุ่มแตกต่างกัน ($H_a : \sigma_A^2 \neq \sigma_B^2$)

2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. สถิติที่ใช้ในการทดสอบคือ F-test การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Classical tests of hypotheses--> Two-group variance-comparison test**

โปรแกรมจะแสดงผลพีอาร์ที่ได้นบนหน้าต่าง Stata Results ดังนี้

```
. sdtest age, by( SA )
```

Variance ratio test

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
0	268	43.15004	.8926391	14.61313	41.39253	44.90755
1	569	52.58904	.5503131	13.12702	51.50814	53.66994
combined	837	49.56676	.4945282	14.30717	48.59609	50.53742

Ho: sd(0) = sd(1)

F(267,568) observed = F_obs = 1.239
F(267,568) lower tail = F_L = 1/F_obs = 0.807
F(267,568) upper tail = F_U = F_obs = 1.239

Ha: sd(0) < sd(1) Ha: sd(0) != sd(1) Ha: sd(0) > sd(1)
P < F_obs = 0.9812 P < F_L + P > F_U = 0.0414 P > F_obs = 0.0188

4. ผลการทดสอบความแตกต่างความแปรปรวนระหว่างตัวอย่าง 2 กลุ่มพบว่า p-value มีค่าเท่ากับ 0.0414 ซึ่ง มีค่าน้อยกว่า 0.05 แสดงว่าความแปรปรวนระหว่างตัวอย่าง 2 กลุ่มแตกต่างกันอย่างมีนัยสำคัญทางสถิติ

ขั้นตอนที่ 2 ทดสอบความแตกต่างของค่าเฉลี่ยระหว่างตัวอย่าง 2 กลุ่ม

- ตั้งสมมติฐานนัยและสมมติฐานทางเลือกดังนี้
สมมติฐานนัยคือ ค่าเฉลี่ยระหว่างประชากร 2 กลุ่ม ไม่แตกต่างกัน ($H_0 : \mu_A = \mu_B$)
สมมติฐานทางเลือกคือ ค่าเฉลี่ยระหว่างประชากร 2 กลุ่มแตกต่างกัน ($H_a : \mu_A \neq \mu_B$)
- กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
- สถิติที่ใช้ในการทดสอบคือ t-test แบบ separated variance การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Classical tests of hypotheses--> Two-group mean-comparison test**

โปรแกรมจะแสดงผลพีอาร์ที่ได้นบนหน้าต่าง Stata Results ดังนี้

```
. ttest age, by(SA) unequal
```

Two-sample t test with unequal variances

Group	Obs	Mean	Std. Err.	Std. Dev.	[95% Conf. Interval]	
0	268	43.15004	.8926391	14.61313	41.39253	44.90755
1	569	52.58904	.5503131	13.12702	51.50814	53.66994
combined	837	49.56676	.4945282	14.30717	48.59609	50.53742
diff		-9.439	1.048642		-11.49954	-7.378464

Satterthwaite's degrees of freedom: 476.193

Ho: mean(0) - mean(1) = diff = 0

Ha: diff < 0 Ha: diff != 0 Ha: diff > 0
t = -9.0012 t = -9.0012 t = -9.0012
P < t = 0.0000 P > |t| = 0.0000 P > t = 1.0000

4. ผลการทดสอบความแตกต่างค่าเฉลี่ยอายุระหว่างกลุ่มผู้ป่วยที่เป็น SA กับกลุ่มผู้ป่วยที่ไม่เป็น SA พบว่า p-value มีค่า <0.001 แสดงว่าค่าเฉลี่ยอายุระหว่างกลุ่มผู้ป่วยทั้ง 2 กลุ่มแตกต่างกันอย่างมีนัยสำคัญทางสถิติ โดยกลุ่มผู้ป่วยที่เป็น SA มีอายุเฉลี่ยเท่ากับ 52.59 ปี ส่วนกลุ่มผู้ป่วยที่ไม่เป็น SA มีอายุเฉลี่ยเท่ากับ 43.15 ปี

2. ลักษณะการกระจายของข้อมูลไม่ใช่โค้งปกติ (Non-normal distribution)

ในกรณีที่ข้อมูลทั้งสองกลุ่มไม่ได้มีการแจกแจงแบบโค้งปกติ เราจะทำการทดสอบความแตกต่างของค่ามัธยฐานระหว่างประชากรสองกลุ่มแทนค่าเฉลี่ย โดยสถิติที่ใช้ในการทดสอบคือ **Wilcoxon rank-sum (Mann-Whitney) test** ซึ่งเป็นสถิติที่ไม่ได้คำนึงถึงการกระจายของข้อมูล ซึ่งเราเรียกวิธีการวิเคราะห์ข้อมูลแบบนี้ว่า Non-parametric ซึ่งเป็นวิธีที่เหมาะสมสำหรับข้อมูลที่มีความเบ้มาก ๆ

ตัวอย่างที่ 4 จากการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิด SA ของจุลชีพและคณะ ในตัวอย่างที่ 3 โดยในตัวอย่างนี้ต้องการทดสอบความแตกต่างของค่า Epworth Sleepiness Scale (ESS) score ระหว่างกลุ่มผู้ป่วยที่เป็น SA กับกลุ่มผู้ป่วยที่ไม่เป็น SA ซึ่งมีขั้นตอนในการทดสอบสมมติฐานดังนี้

ขั้นตอนแรก คือการทดสอบการกระจายของข้อมูลของทั้งสองกลุ่ม (กลุ่มผู้ป่วยที่เป็น SA และกลุ่มผู้ป่วยที่ไม่ได้เป็น SA) ว่ามีการกระจายแบบโค้งปกติหรือไม่ โดยมีขั้นตอนในการทดสอบดังนี้

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือก
สมมติฐานนัย คือข้อมูลมีการกระจายแบบโค้งปกติ
สมมติฐานทางเลือก คือข้อมูลไม่ได้มีการกระจายแบบโค้งปกติ
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. ใช้สถิติ Shapiro-wilk test ในการทดสอบการกระจายของข้อมูลแต่ละกลุ่ม การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Distributional plots and tests --> Shapiro-Wilk normality test**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. by SA, sort : swilk ess

-----
-> SA = 0

                Shapiro-Wilk W test for normal data

    Variable |      Obs      W          V          z        Prob>z
-----+-----
          ess |      268   0.98753     2.405     2.049     0.02025

-----
-> SA = 1

                Shapiro-Wilk W test for normal data

    Variable |      Obs      W          V          z        Prob>z
-----+-----
          ess |      569   0.98837     4.397     3.581     0.00017
```

4. ผลการทดสอบการกระจายของข้อมูลอายุในกลุ่มผู้ป่วยที่เป็น SA และกลุ่มผู้ป่วยที่ไม่ได้เป็น SA พบว่า p-value มีค่าน้อยกว่า 0.05 ทั้งสองกลุ่ม แสดงว่า การกระจายของข้อมูลอายุของทั้งสองกลุ่มไม่ได้มีการแจกแจงแบบโค้งปกติ
5. เมื่อพิจารณาค่าสัมประสิทธิ์ต่างๆ เช่น ค่าความโด่ง (kurtosis) ค่าความเบ้ (skewness) หรือค่าส่วนเบี่ยงเบนมาตรฐาน พบว่า ส่วนเบี่ยงเบนมาตรฐานของข้อมูลทั้งสองกลุ่มมีค่ามากกว่าครึ่งหนึ่งของค่าเฉลี่ย ซึ่งมีผลทำให้ค่าเฉลี่ยลบสองเท่าของส่วนเบี่ยงเบนมาตรฐานเป็นค่าติดลบ ซึ่งเป็นค่าที่เป็นไปไม่ได้ สรุปว่าการกระจายของข้อมูล ESS score ของทั้งสองกลุ่มไม่ได้แจกแจงแบบโค้งปกติ เราจึงใช้สถิติ Mann-Whitney test ในการทดสอบความแตกต่างของค่ามัธยฐานของ ESS score ระหว่างกลุ่มผู้ป่วยที่เป็น SA และกลุ่มผู้ป่วยที่ไม่ได้เป็น SA ในขั้นตอนต่อไป

```
. by SA, sort: sum ess,detail
```

-> SA = 0

ess				
Percentiles		Smallest		
1%	0	0		
5%	1	0		
10%	3	0	Obs	268
25%	6	0	Sum of Wgt.	268
50%	10		Mean	10.16418
		Largest	Std. Dev.	5.665711
75%	14	23		
90%	17	23	Variance	32.10029
95%	19	24	Skewness	.1156905
99%	23	24	Kurtosis	2.382658

-> SA = 1

ess				
Percentiles		Smallest		
1%	0	0		
5%	2	0		
10%	4	0	Obs	569
25%	7	0	Sum of Wgt.	569
50%	11		Mean	10.94376
		Largest	Std. Dev.	5.773635
75%	15	24		
90%	19	24	Variance	33.33486
95%	21	24	Skewness	.0971281
99%	23	24	Kurtosis	2.22615

ขั้นตอนการทดสอบความแตกต่างของค่ามัธยฐานระหว่างกลุ่มตัวอย่าง 2 กลุ่ม

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือกดังนี้
สมมติฐานนัยคือ ค่ามัธยฐานระหว่างประชากร 2 กลุ่ม ไม่แตกต่างกัน ($H_0 : M_A = M_B$)
สมมติฐานทางเลือกคือ ค่ามัธยฐานระหว่างประชากร 2 กลุ่มแตกต่างกัน ($H_a : M_A \neq M_B$)
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. สถิติที่ใช้ในการทดสอบคือ Mann-Whitney test การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Nonparametric tests of hypotheses --> Wilcoxon rank-sum test**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. ranksum ess,by(SA)

Two-sample Wilcoxon rank-sum (Mann-Whitney) test

-----+-----
      SA |      obs   rank sum   expected
-----+-----
       0 |      268    106808    112292
       1 |      569    243895    238411
-----+-----
 combined |      837    350703    350703

unadjusted variance      10649025
adjustment for ties     -27615.655
-----
adjusted variance      10621409

Ho: ess(SA==0) = ess(SA==1)
      z =  -1.683
      Prob > |z| =   0.0924
```

4. ผลการทดสอบความแตกต่างค่ามัธยฐานของ ESS score ระหว่างผู้ป่วยที่เป็น SA กับผู้ป่วยที่ไม่เป็น SA พบว่า p-value มีค่าเท่ากับ 0.0924 แสดงว่าค่ามัธยฐานของ ESS score ระหว่างผู้ป่วย 2 กลุ่มไม่แตกต่างกัน โดยกลุ่มผู้ป่วยที่เป็น SA มีค่ามัธยฐานเท่ากับ 11 และมีค่าต่ำสุด-สูงสุดเท่ากับ 0-24 ส่วนในกลุ่มผู้ป่วยที่ไม่เป็น SA มีค่ามัธยฐานเท่ากับ 10 และมีค่าต่ำสุด-สูงสุดเท่ากับ 0-24

บทที่ 5

การทดสอบสมมติฐานสำหรับข้อมูลกลุ่ม (Hypothesis testing for categorical data)

การทดสอบสมมติฐานสำหรับข้อมูลกลุ่มเป็นการทดสอบความแตกต่างของสัดส่วนของเหตุการณ์ที่สนใจในประชากร หรือเป็นการหาความสัมพันธ์ระหว่างตัวแปรกลุ่มสองตัว โดยในบทนี้จะกล่าวถึงการทดสอบสมมติฐานในประชากรสองกลุ่ม และการคำนวณหาขนาดของความสัมพันธ์ระหว่างตัวแปรกลุ่มสองตัวเท่านั้น ซึ่งเป็นการทดสอบที่พบได้บ่อยในงานวิจัยทางการแพทย์

การทดสอบสมมติฐานในประชากรสองกลุ่มที่อิสระต่อกัน

สถิติที่ใช้ในการทดสอบความแตกต่างของสัดส่วนของเหตุการณ์ที่สนใจระหว่างประชากรสองกลุ่มที่อิสระต่อกันคือ **Chi-square test** หรือ **Fisher Exact test** ซึ่งมีข้อตกลงสำหรับการใช้สถิติ Chi-square คือ ค่าคาดหวัง (Expected value) ของแต่ละเซลล์จะต้องมีค่ามากกว่า 5 หรือมีจำนวนเซลล์ที่มีค่าคาดหวังน้อยกว่า 5 ได้ไม่เกิน 20% ของจำนวนเซลล์ทั้งหมด โดยที่ค่าคาดหวังของเซลล์นั้นสามารถคำนวณได้จากผลคูณของผลรวมในแนวนอน และในแนวคอลัมน์ที่เซลล์นั้นอยู่ หาดด้วยจำนวนตัวอย่างทั้งหมด แต่ถ้าไม่เป็นไปตามข้อตกลงดังกล่าว ต้องใช้สถิติ Fisher Exact test ในการทดสอบสมมติฐานแทน

ตัวอย่างที่ 1 จากการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิด sleep apnea (SA) โดยจุลวิทย์และคณะ เพื่อนำไปใช้ในการสร้างคะแนนสำหรับทำนายการเกิด SA ซึ่งมีรูปแบบการศึกษาแบบภาคตัดขวาง (cross-sectional study) ผู้วิจัยต้องการทราบว่าสัดส่วนการเกิด stop breathing ระหว่างกลุ่มผู้ป่วยที่เป็น SA กับกลุ่มผู้ป่วยที่ไม่เป็น SA มีความแตกต่างกันหรือไม่

ขั้นตอนการทดสอบสมมติฐาน

1. ตั้งสมมติฐานนัยและสมมติฐานทางเลือกดังนี้
สมมติฐานนัยคือ สัดส่วนการเกิด stop breathing ระหว่างผู้ป่วย 2 กลุ่ม ไม่แตกต่างกัน ($H_0 : P_1 = P_2$)
สมมติฐานทางเลือกคือ สัดส่วนการเกิด stop breathing ระหว่างผู้ป่วย 2 กลุ่มแตกต่างกัน ($H_a : P_1 \neq P_2$)
2. กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
3. สถิติที่ใช้ในการทดสอบคือ Chi-square test หรือ Fisher Exact test โดยขึ้นอยู่กับค่าคาดหวังของแต่ละเซลล์ การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Tables --> Two-way tables with measures of association**

โปรแกรมจะแสดงผลที่ได้บนหน้าต่าง Stata Results ดังนี้

```
. tab stop_bre SA,col exp chi2 exact
```

Key	frequency	expected frequency	column percentage
stop breathing	SA		
	0	1	Total
yes	93	396	489
	156.6	332.4	489.0
	34.70	69.60	58.42
no	175	173	348
	111.4	236.6	348.0
	65.30	30.40	41.58
Total	268	569	837
	268.0	569.0	837.0
	100.00	100.00	100.00
Pearson chi2(1) = 91.3257 Pr = 0.000			
Fisher's exact = 0.000			
1-sided Fisher's exact = 0.000			

4. พิจารณาค่าคาดหวังในแต่ละเซลล์ (แถวที่ 2) พบว่ามีค่ามากกว่า 5 ทุกเซลล์ เพราะฉะนั้นสถิติที่เหมาะสมสำหรับการทดสอบสมมติฐานในกรณีนี้คือ Chi-square test ซึ่งจากผลการทดสอบความแตกต่างของสัดส่วนของการเกิด stop breathing ระหว่างกลุ่มผู้ป่วยที่เป็น sleep apnea กับผู้ป่วยที่ไม่เป็น sleep apnea พบว่า p-value มีค่า <0.001 แสดงว่าสัดส่วนของการเกิด stop breathing ระหว่างผู้ป่วย 2 กลุ่มแตกต่างกันอย่างมีนัยสำคัญทางสถิติ โดยพบว่าสัดส่วนของการเกิด stop breathing ในกลุ่มผู้ป่วยที่เป็น sleep apnea เท่ากับ 69.6% ในขณะที่สัดส่วนของการเกิด stop breathing ในกลุ่มผู้ป่วยที่ไม่เป็น sleep apnea เท่ากับ 34.7% นั้นหมายความว่า การเกิด stop breathing มีความสัมพันธ์กับการเป็น sleep apnea

ตัวอย่างที่ 2 จากการศึกษาเปรียบเทียบการตอบสนองต่อเชื้อไวรัสของยา Efavirenz (EFV) และ Nevirapine (NVP) ในผู้ป่วยที่ติดเชื้อเอชไอวี โดยวิธีสุ่มและคณะ ซึ่งมีรูปแบบการศึกษาแบบ Randomized control trial ผู้วิจัยต้องการทราบว่า อัตราการตอบสนองต่อเชื้อไวรัสระหว่างผู้ป่วยที่ได้รับยา EFV และ NVP มีความแตกต่างกันหรือไม่

ขั้นตอนการทดสอบสมมติฐาน

- ตั้งสมมติฐานนัยและสมมติฐานทางเลือกดังนี้
สมมติฐานนัยคือ สัดส่วนของผู้ป่วยที่ตอบสนองต่อเชื้อไวรัส ระหว่างกลุ่มที่ได้รับยา EFV และ NVP ไม่แตกต่างกัน ($H_0 : P_1 = P_2$)
สมมติฐานทางเลือกคือ สัดส่วนของผู้ป่วยที่ตอบสนองต่อเชื้อไวรัส ระหว่างกลุ่มที่ได้รับยา EFV และ NVP แตกต่างกัน ($H_a : P_1 \neq P_2$)
- กำหนดระดับนัยสำคัญของการทดสอบสมมติฐานเท่ากับ 0.05
- สถิติที่ใช้ในการทดสอบคือ Chi-square test หรือ Fisher Exact test โดยขึ้นอยู่กับค่าคาดหวังของแต่ละเซลล์ การวิเคราะห์ข้อมูลโดยใช้โปรแกรม STATA สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Summaries, tables, & tests --> Tables --> Two-way tables with measures of association**

เมื่อ I_{E^+} คือ อุบัติการณ์ของการเกิดโรคหรือเหตุการณ์ที่สนใจในกลุ่มที่ได้รับปัจจัยเสี่ยง

I_{E^-} คือ อุบัติการณ์ของการเกิดโรคหรือเหตุการณ์ที่สนใจในกลุ่มที่ไม่ได้รับปัจจัยเสี่ยง

ซึ่งค่า RR เหมาะสำหรับการแสดงขนาดความสัมพันธ์ในกรณีที่รูปแบบการวิจัยเป็นแบบ Cohort study เท่านั้น

ตัวอย่างที่ 3 ในการศึกษาอัตราการรอดชีวิตของไตในผู้ป่วยที่ได้รับการผ่าตัดเปลี่ยนไตที่โรงพยาบาลรามาริบัติ โดย อัมรินทร์และคณะ (วารสารวิทยาการระบาด 2541) ผู้วิจัยต้องการทราบว่าผู้ป่วยที่ได้รับเปลี่ยนไตที่ได้จากผู้มีชีวิต (Living) กับได้รับเปลี่ยนไตที่ได้จากผู้เสียชีวิต (Cadavor) มีความสัมพันธ์กับการเกิดภาวะไตล้มเหลวหรือไม่ โดยผู้วิจัยได้ทำการศึกษาในผู้ป่วยเปลี่ยนไตจำนวน 335 คน ในจำนวนนี้มีผู้ป่วยที่ได้รับไตจากผู้มีชีวิต 128 คน และได้รับไตจากผู้เสียชีวิต 207 คน โดยพบว่าในจำนวนผู้ป่วยที่ได้รับไตจากผู้มีชีวิต เกิดภาวะไตล้มเหลวจำนวน 5 คน และในจำนวนผู้ป่วยที่ได้รับไตจากผู้เสียชีวิต เกิดภาวะไตล้มเหลวจำนวน 74 คน การหาขนาดความสัมพันธ์โดยใช้โปรแกรม STATA สำหรับข้อมูลที่มีการนับจำนวนของเหตุการณ์ที่สนใจแล้ว สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Epidemiology and related --> Tables for epidemiology --> Cohort study risk-ratio etc. calculator**

```
. csi 74 5 133 123
```

	Exposed	Unexposed	Total
Cases	74	5	79
Noncases	133	123	256
Total	207	128	335
Risk	.3574879	.0390625	.2358209
	Point estimate		[95% Conf. Interval]
Risk difference	.3184254		.2450152 .3918357
Risk ratio	9.151691		3.80194 22.02913
Attr. frac. ex.	.8907306		.7369764 .9546056
Attr. frac. pop	.8343552		

	chi2(1) = 44.50		Pr>chi2 = 0.0000

จากผลลัพธ์ที่ได้พบว่า RR ของการศึกษานี้มีค่าเท่ากับ 9.15 นั่นคือ ผู้ป่วยที่ได้รับไตจากผู้เสียชีวิตมีโอกาสที่จะเกิดภาวะไตล้มเหลวประมาณ 9 เท่า เมื่อเทียบกับผู้ป่วยที่ได้รับไตจากผู้มีชีวิต โดยมี 95% CI of RR อยู่ระหว่าง 3.80-22.03

2. Odds ratio (OR)

ในกรณีที่รูปแบบการศึกษาเป็นแบบ Case-control หรือ Cross-sectional study เราไม่สามารถคำนวณหาอุบัติการณ์ของการเกิดโรคหรือเหตุการณ์ที่สนใจได้ ดังนั้นค่า OR จึงถูกนำมาใช้เป็นตัวแทนของขนาดความสัมพันธ์ในกรณีนี้

การคำนวณค่า OR สามารถทำได้ดังนี้

$$OR = \frac{Odds_{E^+}}{Odds_{E^-}}$$

เมื่อ $Odds = \frac{P}{1-P}$ ซึ่งก็คือ สัดส่วนของความน่าจะเป็นของการเป็นโรคต่อความน่าจะเป็นของการไม่เป็นโรค

ตัวอย่างที่ 4 ในการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิดกระดูกข้อสะโพกหัก (Hip fracture) โดยไพลุลย์และคณะ ผู้วิจัยต้องการทราบว่าคนที่ผู้ป่วยมีประวัติการรับประทุษร้ายแผนโบราณเสี่ยงต่อการเกิดกระดูกข้อสะโพกหักหรือไม่ โดยผู้วิจัยได้ทำการศึกษาในผู้ป่วยจำนวน 452 คน ในจำนวนนี้มีผู้ป่วยที่มีประวัติการรับประทุษร้ายแผนโบราณจำนวน 28 คน และผู้ป่วยที่ไม่มีประวัติการรับประทุษร้ายแผนโบราณจำนวน 424 คน โดยพบว่าในจำนวนผู้ป่วยที่มีประวัติการรับประทุษร้ายแผนโบราณ เกิดกระดูกข้อสะโพกหักจำนวน 20 คน และในจำนวนผู้ป่วยที่ไม่มีประวัติการรับประทุษร้ายแผนโบราณ เกิดกระดูกข้อสะโพกหักจำนวน 216 คน การหาขนาดความสัมพันธ์โดยใช้โปรแกรม STATA สำหรับข้อมูลที่มีการนับจำนวนของเหตุการณ์ที่สนใจแล้ว สามารถทำได้ดังนี้

เลือกเมนู **Statistics --> Epidemiology and related --> Tables for epidemiology --> Case-control odds-ratio calculator**

```

. cci 20 216 8 208

              |      Exposed      Unexposed      |      Total      Proportion
              +-----+-----+-----+-----+
              |      Cases      |      20      216      |      236      0.0847
              |      Controls      |      8      208      |      216      0.0370
              +-----+-----+-----+-----+
              |      Total      |      28      424      |      452      0.0619
              |      Point estimate      |      [95% Conf. Interval]
              +-----+-----+-----+-----+
              |      Odds ratio      |      2.407407      |      .9872681      6.451643 (exact)
              |      Attr. frac. ex. |      .5846154      |      -.0128961      .8450007 (exact)
              |      Attr. frac. pop |      .0495437      |
              +-----+-----+-----+-----+
              |      chi2(1) =      4.42      Pr>chi2 = 0.0356

```

จากผลลัพธ์ที่ได้พบว่า OR ของการศึกษานี้มีค่าเท่ากับ 2.41 นั่นคือ ผู้ป่วยเกิดกระดูกสะโพกหักมีโอกาสที่จะมีประวัติการรับประทานยาแผนโบราณสูงกว่า 2.4 เท่า เมื่อเทียบกับผู้ป่วยที่ไม่เกิดกระดูกหัก โดยมี 95% CI of OR อยู่ระหว่าง 0.99-6.45

บทที่ 6

การนำเสนอผลการวิจัย (Reporting results)

ในการทำงานวิจัย ผู้วิจัยควรออกแบบตารางการนำเสนอผลการวิจัย หรือเรียกว่าตารางหุ่น (Dummy tables) ตั้งแต่ขั้นตอนการวางแผนโครงร่างวิจัย ทั้งนี้การออกแบบตารางหุ่นจะขึ้นอยู่กับรูปแบบงานวิจัยและวัตถุประสงค์ของงานวิจัยเป็นหลัก โดยการนำเสนอผลการวิจัยจะเริ่มจากตารางสำหรับการบรรยายลักษณะของตัวอย่าง (Description of patients' s characteristics) ก่อนเสมอ ตามด้วยตารางสำหรับการวิเคราะห์แบบตัวแปรเดียว (Univariate analysis) และตารางสำหรับการวิเคราะห์แบบพหุ (Multivairate analysis) ตามตัวอย่างดังต่อไปนี้

ตัวอย่างที่ 1 การนำเสนอผลการวิจัยเรื่องการศึกษาปัจจัยที่มีความสัมพันธ์กับการเกิด Sleep apnea ของ Rodsutti et.al. ซึ่งปัจจัยที่สนใจได้แก่ อายุ เพศ ดัชนีความอ้วน การสูบบุหรี่ การดื่มสุรา การมีประวัติการนอนกรน การมีประวัติการหยุดหายใจขณะนอนหลับ โดยมีรูปแบบการวิจัยแบบภาคตัดขวาง (Cross-sectional study) มีขั้นตอนการนำเสนอผลการวิจัยดังนี้

- ตารางที่ 1 บรรยายลักษณะทั่วไปของกลุ่มตัวอย่างทั้งหมด โดยใช้หลักการของสถิติเชิงพรรณนา (ดังรายละเอียดในบทที่ 2)
- ตารางที่ 2 ทดสอบความแตกต่างของปัจจัยต่างๆระหว่างกลุ่มที่เป็น sleep apnea กับกลุ่มที่ไม่ได้เป็น sleep apnea ซึ่งเป็นการทดสอบความสัมพันธ์ของปัจจัยที่ละตัวกับการเกิด sleep apnea หรือที่เรียกว่า การวิเคราะห์แบบตัวแปรเดียว (Univariate analysis) โดยใช้หลักการของการทดสอบสมมติฐาน (ดังรายละเอียดในบทที่ 4 และ 5)
- ตารางที่ 3 วิเคราะห์หาความสัมพันธ์ระหว่างปัจจัยหลายๆตัวกับการเกิด sleep apnea หรือที่เรียกว่า การวิเคราะห์แบบพหุ (Multivairate analysis) ซึ่งไม่ได้กล่าวรายละเอียดการวิเคราะห์ไว้ในหนังสือเล่มนี้

Table 1 Description of patients's characteristics

Characteristics	N (%)
Age, mean(sd)	
Age, n(%)	
<30	
30-44	
45-59	
≥ 60	
Gender	
Male	
Female	
BMI, mean(sd)	
BMI	
<25	
25-29.9	
30-39.9	
≥ 40	
Snoring	
Yes	
No	
Stopping breathing	
Yes	
No	
Waking up refreshed	
Yes	
Sometimes	
No	
Leg kicking	
Yes	
Sometimes	
No	
ESS score, median(range)	

Table 2 Description of patients's characteristics between SA and non-SA

Characteristics	SA n(%)	Non-SA n(%)	P value
Age, mean(sd)			
Age, n(%)			
<30			
30-44			
45-59			
≥ 60			
Gender			
male			
female			
BMI, mean(sd)			
BMI			
<25			
25-29.9			
30-39.9			
≥ 40			
Snoring			
Yes			
No			
Stopping breathing			
Yes			
No			
Waking up refreshed			
Yes			
Sometimes			
No			
Leg kicking			
Yes			
Sometimes			
No			
ESS score, median(range)			

Table3 Factors associated with SA (AHI $\geq$ 5): multiple logistic regression analysis

Factors	Coefficient	SE	P value	OR (95% CI)
Age				
$\geq$ 60				
45-59				
30-44				
< 30				
Gender				
Male				
Female				
Snoring				
Yes				
No				
Stopping breathing				
Yes				
No				
BMI				
$\geq$ 40				
30-39.9				
25-29.9				
< 25				

ตัวอย่างที่ 2 การนำเสนอผลการวิจัยเรื่องการเปลี่ยนแปลงของความหนาแน่นกระดูก (BMD) และดัชนีทางชีวเคมีของการสร้างและสลายมวลกระดูก (Biochemical markers of bone remodeling) ซึ่งได้แก่ serum C-termina-telopeptide of type I collagen (CTx), PINP, PTH และ Vitamin D ภายหลังจากได้รับแคลเซียมเสริม โดยใช้หลักการของการทดสอบสมมติฐานสำหรับข้อมูลต่อเนื่องในประชากรสองกลุ่มที่ไม่อิสระต่อกัน (ดังรายละเอียดในบทที่ 4)

Table 4 Describe BMD and biochemical markers between randomization and 96 weeks

Parameters	Randomization	96 weeks	P value
BMD			
L2L4			
Neck			
Wards			
Trochanter			
Shaft			
Biochemical markers			
CTx			
PINP			
PTH			
Vitamin D			

References

1. Daniel WW. Biostatistics: A foundation for analysis in the health sciences, sixth edition. New York: John Wiley & Sons 1995.
2. Michael JC, David M. Medical Statistics, second edition. New York: John Wiley & Sons 1993.
3. Lang TA, Secic M. How to report statistics in medicine: Annotated guidelines for authors, editors, and reviewers. Philadelphia: American College of Physician 1997.
4. Rodsutti J, Hensley M, Thakkestian A, D'Este C, Attia J. A clinical decision rule to prioritize polysomnography in patients with suspected sleep apnea. Sleep. 2004 Jun 15; 27(4): 694-9.
5. อัมรินทร์ ทักขิณเสถียร. สถิติเบื้องต้นสำหรับงานวิจัยทางการแพทย์. เอกสารประกอบการสอนการอบรมเชิงปฏิบัติการระบาดวิทยาคลินิกพื้นฐานและสถิติเบื้องต้น สำหรับแพทย์ผู้ช่วยอาจารย์ และอาจารย์แพทย์. กลุ่มงานระบาดวิทยาคลินิกและชีวสถิติ สำนักงานวิจัยคณะแพทยศาสตร์โรงพยาบาลรามาธิบดี