

บทที่ 3 ความรู้สถิติพื้นฐาน

ความหมายของสถิติ

สถิติ (statistics) อาจให้ความหมายกว้างๆ ได้ 2 ประการ คือ

1. สถิติ หมายถึง ข้อความจริง หรือตัวเลขซึ่งได้รวบรวมไว้เพื่อความหมายที่แน่นอน เช่น สถิติพลเมือง สถิติจำนวนอุบัติเหตุในรอบปี สถิติจำนวนผู้ป่วยที่มารับการรักษายาบาลในโรงพยาบาลหนึ่ง สถิติคนไข้เป็นโรคมะเร็ง ฯลฯ เป็นต้น

2. สถิติ หมายถึง ศาสตร์ที่ว่าด้วยวิธีการเก็บรวบรวมข้อความจริงและตัวเลขที่แสดงข้อเท็จจริงซึ่งเรียกว่า ข้อมูล การนำเสนอข้อมูล การวิเคราะห์ข้อมูล และการตีความตลอดจนการสรุปผลข้อมูล เพื่อใช้เป็นประโยชน์ในการตัดสินใจที่มีเหตุผล

สถิติในความหมายแรก จะหมายถึงสถิติในฐานะที่เป็นตัวเลขซึ่งเรียกว่า ข้อมูลทางสถิติ ส่วนสถิติในความหมายที่สองจะเป็นศาสตร์ที่เรียกว่า สถิติศาสตร์ ซึ่งเป็นศาสตร์แขนงหนึ่งที่ประยุกต์มาจากคณิตศาสตร์ ในส่วนของทฤษฎีทางสถิติจึงมีความเกี่ยวข้องกับคณิตศาสตร์อย่างมาก เช่น ทฤษฎีการแจกแจงความน่าจะเป็น ทฤษฎีการประมาณ ทฤษฎีการทดสอบสมมติฐาน ในทางปฏิบัติ กระบวนการทางสถิติ (ในความหมายที่สอง) จะต้องดำเนินการตามระเบียบวิธีทางสถิติ ได้แก่ การเก็บรวบรวมข้อมูล การนำเสนอข้อมูล การวิเคราะห์ข้อมูล และการแปลความหมายหรือการตีความข้อมูล

ระเบียบวิธีการทางสถิติ

ระเบียบวิธีการทางสถิติ ประกอบด้วย

1. การเก็บรวบรวมข้อมูล (data collection)

เป็นการรวบรวมข้อมูลจากแหล่งข้อมูลตามที่ได้มีการวางแผนไว้ ซึ่งอาจเป็นได้ทั้งข้อมูลปฐมภูมิ หรือทุติยภูมิ

2. การนำเสนอข้อมูล (data presentation)

เป็นการจัดทำข้อมูลที่รวบรวมได้ให้อยู่ในรูปแบบที่กะทัดรัด เช่น ตาราง กราฟ แผนภูมิ ข้อความ เป็นต้น เพื่อความสะดวกในการอ่านข้อมูล ให้เข้าใจง่าย และเพื่อประโยชน์ในการวิเคราะห์ต่อไป

3. การวิเคราะห์ข้อมูล (data analysis)

เป็นขั้นตอนการประมวลผลข้อมูล ซึ่งในการวิเคราะห์จำเป็นต้องใช้สูตรทางสถิติต่างๆ หรือใช้อ้างอิงทางสถิติ ขึ้นกับวัตถุประสงค์ของงานนั้นๆ เช่น การวิเคราะห์แนวโน้มเข้าสู่ส่วนกลาง การวัดการกระจาย การทดสอบสมมติฐาน การประมาณค่า เป็นต้น

4. การแปลความหมาย (interpretation)

เป็นขั้นตอนของการนำผลการวิเคราะห์มาอธิบายให้บุคคลทั่วไปเข้าใจ อาจจำเป็นต้องมีการขยายความในการอธิบาย เพื่อให้งานที่ศึกษาเป็นประโยชน์ต่อคนทั่วไปได้

จากกระบวนการทางสถิติดังกล่าว เราสามารถจำแนกเป็นสถิติศาสตร์ ที่สอดคล้องกับขั้นตอนต่างๆ ได้ 2 ลักษณะคือ สถิติบรรยาย (หรือสถิติเชิงพรรณนา) และสถิติอ้างอิง (หรือสถิติเชิงอนุมาน)

สถิติบรรยาย และสถิติอ้างอิง

สถิติบรรยาย หรือสถิติเชิงพรรณนา (descriptive statistics)

เป็นการอธิบายลักษณะของข้อมูลในรูปของการบรรยายลักษณะต่างๆ ไปของข้อมูล โดยจัดนำเสนอเป็นบทความ บทความกิ่งตาราง แสดงด้วยกราฟ หรือแผนภูมิ ตลอดจนทำเป็นรูปภาพต่างๆ มีการคำนวณหาความหมายของข้อมูลโดยวิธีทางสถิติอย่างง่าย เพื่อให้เป็นรูปแบบของข้อมูลในเบื้องต้นให้สามารถตีความหมายของข้อมูลได้ตามความจริง

สถิติบรรยายนี้อาจทำการศึกษาเกี่ยวกับข้อมูลที่เป็นกลุ่มเล็กๆ หรือกลุ่มใหญ่ โดยทั่วไปก็ได้ และผลการวิเคราะห์จะใช้อธิบายเฉพาะกลุ่มที่นำมาศึกษาเท่านั้น

สถิติบรรยายที่ใช้ในงานวิจัย เช่น การแจกแจงความถี่ ร้อยละ การวัดแนวโน้มเข้าสู่ส่วนกลาง การวัดการกระจาย เป็นต้น

สถิติอ้างอิง หรือสถิติเชิงอนุมาน (inferential statistics)

เป็นเทคนิคที่นำข้อมูลเพียงส่วนหนึ่งไปอธิบายเกี่ยวกับข้อมูลส่วนใหญ่โดยทั่วไป โดยใช้พื้นฐานเรื่องความน่าจะเป็น เป็นหลักในการอนุมาน หรือทำนายไปยังกลุ่มประชากรเป้าหมาย การใช้สถิติอ้างอิงทำได้ 2 ลักษณะ คือ การประมาณค่าประชากร และการทดสอบสมมติฐาน

เพื่อให้มองเห็นข้อแตกต่างระหว่างสถิติบรรยาย และสถิติอ้างอิง และมองเห็นลักษณะของสถิติอ้างอิงได้อย่างเด่นชัดขึ้น จะขออธิบายความหมายของคำที่เกี่ยวข้องต่อไปนี้

ประชากร (population) หมายถึง ขอบเขตของข้อมูลทั้งหมดที่เรากำลังทำการศึกษา หรืออาจหมายถึง กลุ่มของสิ่งของทั้งหมดที่ให้ข้อมูลตามที่เราต้องการศึกษา เช่น ศึกษาเกี่ยวกับคนไข้สูติ-นรีเวชของโรงพยาบาลมหาราชนครเชียงใหม่ ในปี 2548 ทั้งหมด ซึ่งอาจดูได้จากประวัติผู้ป่วย เป็นต้น

ลักษณะของประชากรที่ศึกษา อาจมีจำนวนจำกัด (finite population) ดังตัวอย่างข้างต้น หรืออาจมีจำนวนอนันต์ (infinite population) เช่น การศึกษาเกี่ยวกับประสิทธิภาพของยาชนิดหนึ่ง ประชากรจะเป็นผลการทดสอบประสิทธิภาพของยาในผู้ป่วยที่ใช้ยานี้ ซึ่งไม่สามารถบอกถึงจำนวนทั้งหมดได้

ค่าที่ประมวลได้จากข้อมูลทั้งหมดของประชากร โดยวิธีการทางสถิติ จะเรียกว่า พารามิเตอร์ (parameter) นิยมใช้สัญลักษณ์อักษรกรีกแทน เช่น

ค่าเฉลี่ยของประชากร แทนด้วย μ อ่านว่า มิว (mu)

ส่วนเบี่ยงเบนมาตรฐานของประชากร แทนด้วย σ อ่านว่า ซิกมา (sigma)

ค่าสัมประสิทธิ์สหสัมพันธ์ของประชากร แทนด้วย ρ อ่านว่า โร (rho)

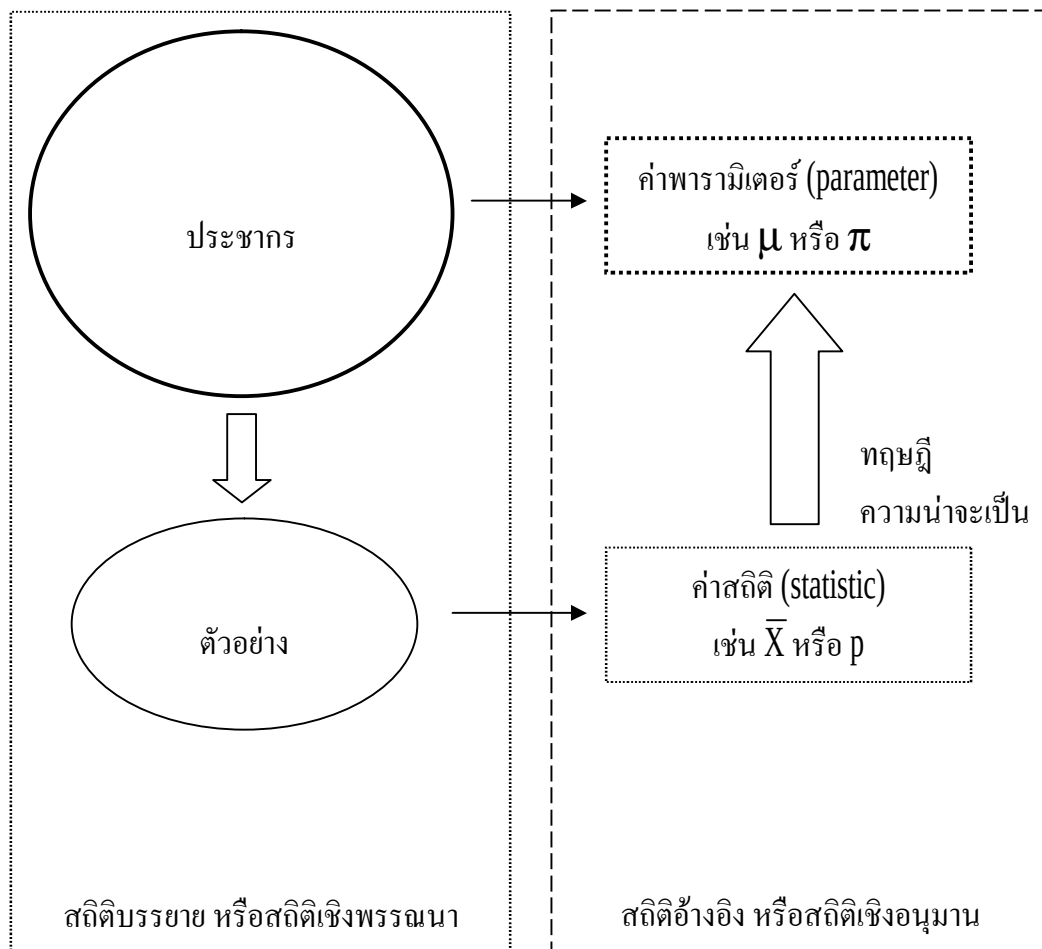
ตัวอย่าง (sample) หมายถึง ส่วนหนึ่งของประชากรซึ่งถูกเลือกมาศึกษา เนื่องจากในบางครั้งพบว่าการศึกษาบางอย่างไม่อาจทำทั้งหมดของประชากรได้ เพราะต้องเสียค่าใช้จ่ายมาก เสียเวลา อาจหาประชากรทั้งหมดไม่ได้ หรือไม่สามารถกระทำกับประชากรทั้งหมดได้ จึงจำเป็นต้องเลือกตัวอย่างมาศึกษา

ค่าที่ประมวลได้จากข้อมูลของตัวอย่าง โดยวิธีการทางสถิติ จะเรียกว่า ค่าสถิติ (statistic) เช่น

ค่าเฉลี่ยของตัวอย่าง แทนด้วย $\bar{X}$

ส่วนเบี่ยงเบนมาตรฐานของตัวอย่าง แทนด้วย S

ค่าสัมประสิทธิ์สหสัมพันธ์ของตัวอย่าง แทนด้วย r



แสดงความสัมพันธ์ระหว่าง

ความจริงที่ได้จากประชากร (พารามิเตอร์) กับผลลัพธ์ที่ได้จากกลุ่มตัวอย่าง (ค่าสถิติ)
กระบวนการทางสถิติที่ใช้ในการอธิบายผลลัพธ์ที่ได้จากกลุ่มตัวอย่าง (สถิติบรรยาย)
และการอนุมานไปสู่การสรุปในระดับประชากร (สถิติอ้างอิง)

สถิติบรรยาย (descriptive statistics)

สถิติบรรยาย หรือสถิติเชิงพรรณนา (descriptive statistics) เป็นสถิติวิเคราะห์เบื้องต้นที่ใช้ในการนำเสนอข้อมูลที่รวบรวมมาจากกลุ่มตัวอย่าง หรือจากประชากร เป็นการบรรยายลักษณะของข้อมูลที่รวบรวม โดยสรุปลักษณะที่สำคัญของข้อมูล ในการนำเสนอสถิติบรรยายประกอบด้วย

การนำเสนอข้อมูล โดยสามารถนำเสนอในรูปของตารางแจกแจงความถี่ อาจเป็นแบบทางเดียวหรือหลายทาง การนำเสนอในรูปแผนภูมิ หรือกราฟ และการนำเสนอในรูปบทความ อธิบาย เป็นต้น

การนำเสนอตัวแทนของข้อมูล ซึ่งเราเรียกว่า การวัดแนวโน้มเข้าสู่ส่วนกลาง โดยทั่วไปจะนำเสนอค่ากลาง 3 ชนิด คือ ค่าเฉลี่ย (mean) ค่ามัธยฐาน (median) และฐานนิยม (mode)

การนำเสนอตำแหน่งหรือลำดับของข้อมูล ประกอบด้วยการแสดงตำแหน่งข้อมูลแบบค่าควอไทล์ (quartile) เดไซล์ (decile) และเปอร์เซนไทล์ (percentile)

การอธิบายการกระจายของข้อมูล ว่ามีการกระจายจากตัวแทนของข้อมูลมากน้อยเพียงใด ได้แก่ ค่าส่วนเบี่ยงเบนเฉลี่ย (mean deviation) ส่วนเบี่ยงเบนมาตรฐาน (standard deviation) พิสัย (range) เป็นต้น

การอธิบายการแจกแจงความถี่ของข้อมูล ในรูปหลายเหลี่ยมความถี่ (frequency polygon) ฮิสโตแกรม (histogram) โค้งความถี่ (frequency curve) ค่าสัมประสิทธิ์ความเบ้ (skewness) สัมประสิทธิ์ความโด่ง (kurtosis) เป็นต้น

การนำเสนอข้อมูล

การแจกแจงความถี่

เป็นการนำเสนอข้อมูลทั้งหมดที่รวบรวมได้ให้อยู่ในรูปแบบที่ชัดเจน เข้าใจง่าย เพื่อให้เกิดความสะดวกในการอ่านข้อมูล และเพื่อประโยชน์ในการวิเคราะห์ข้อมูลต่อไป โดยทั่วไปตารางแจกแจงความถี่จะแบ่งเป็น ตารางแจกแจงความถี่ข้อมูลเชิงคุณภาพ และตารางแจกแจงความถี่ข้อมูลเชิงปริมาณ

ตารางแจกแจงความถี่ข้อมูลเชิงคุณภาพ

ตัวอย่างการนำเสนอตารางแจกแจงความถี่ข้อมูลเชิงคุณภาพ และผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป ดังแสดง

ตารางที่ ... แสดงจำนวนและร้อยละของกลุ่มตัวอย่าง จำแนกตามระดับการศึกษา

ระดับการศึกษา	จำนวน (ราย)	ร้อยละ
ไม่ได้รับการศึกษา	21	15.0
ประถม	109	77.9
มัธยมต้น	6	4.3
มัธยมปลาย	3	2.1
ปริญญาตรี หรือสูงกว่า	1	0.7
รวม	140	100.0

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

Frequencies

Patients education

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid No education	21	15.0	15.0	15.0
Elementary school	109	77.9	77.9	92.9
Middle grade	6	4.3	4.3	97.1
High school	3	2.1	2.1	99.3
College / University	1	.7	.7	100.0
Total	140	100.0	100.0	

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม EPI Info)

EDUCATE	Frequency	Percent	Cum Percent	
0	21	15.0%	15.0%	
1	109	77.9%	92.9%	
2	6	4.3%	97.1%	
3	3	2.1%	99.3%	
4	1	0.7%	100.0%	
Total	140	100.0%	100.0%	

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม MINITAB)

Tally for Discrete Variables: EDUCATE

EDUCATE	Count	Percent
0	21	15.00
1	109	77.86
2	6	4.29
3	3	2.14
4	1	0.71
N=	140	

ตารางแจกแจงความถี่ข้อมูลเชิงปริมาณ

ในการสร้างตารางแจกแจงความถี่ของข้อมูลเชิงปริมาณ กลุ่มของสิ่งสนใจจะเป็นค่าสังเกต ซึ่งเรียกว่า อันตรภาคชั้น (class interval) ในการจัดทำตารางจะมีการคำนวณค่าที่เกี่ยวข้องต่างๆ ได้แก่

พิสัย (range) เป็นความแตกต่างระหว่างค่าสังเกตสูงสุดและค่าสังเกตต่ำสุด

อันตรภาคชั้น (class interval) เป็นความกว้างของชั้นคะแนนของค่าสังเกต (ความแตกต่างของขอบเขตบนและขอบเขตล่าง ของแต่ละอันตรภาคชั้น ความแตกต่างนี้ใช้แทนขนาดของอันตรภาคชั้น)

ขีดจำกัดชั้น (class limit) เป็นตัวเลขเริ่มต้นและลงท้ายของแต่ละอันตรภาคชั้น เลขที่มีค่าน้อยกว่า เรียกว่า ขีดจำกัดล่าง (lower limit) และ เลขที่มีค่ามากกว่า เรียกว่า ขีดจำกัดบน (upper limit)

ขอบเขตชั้น (class boundary) เป็นค่าที่แบ่งแยกอาณาเขตของแต่ละอันตรภาคชั้น หาได้โดยเฉลี่ยขีดจำกัดบน และขีดจำกัดล่าง ของชั้นที่ติดกัน เรียกเลขที่มีค่าน้อยกว่าว่า ขอบเขตล่าง หรือขีดจำกัดล่างจริง (true lower limit) และเลขที่มีค่ามากกว่าว่า ขอบเขตบน หรือขีดจำกัดบนจริง (true upper limit)

จุดกึ่งกลาง (mid point) เป็นค่าเฉลี่ยของขอบเขตบน และขอบเขตล่าง ใช้เป็นตัวแทนของค่าสังเกตต่างๆ ในแต่ละอันตรภาคชั้น

ส่วนการแสดงจำนวนชุดข้อมูล หรือความถี่ของข้อมูล นั้น สามารถแสดงความถี่ในรูปแบบของความถี่สะสม ความถี่สัมพัทธ์ และความถี่สะสมสัมพัทธ์

ความถี่สะสม เป็นการรวมความถี่จากอันตรภาคชั้นที่มีค่าสังเกตน้อยไปยังชั้นที่มีค่าสังเกตมาก หรือเป็นการรวมในทางตรงข้ามก็ได้

ความถี่สัมพัทธ์ เป็นสัดส่วนของความถี่ของอันตรภาคชั้น กับจำนวนค่าสังเกตทั้งหมด โดยทั่วไปนิยมนำเสนอในรูปร้อยละ

ความถี่สะสมสัมพัทธ์ เป็นการสะสมความถี่สัมพัทธ์ โดยมีลักษณะเช่นเดียวกันกับความถี่สะสม

ตัวอย่างการนำเสนอตารางแจกแจงความถี่ข้อมูลเชิงปริมาณ ข้อมูลอายุของผู้ป่วยจำนวน 140 ราย (52, 45, 37, 34, 39, 43, 46, 84, 87, 76) และผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป และการนำเสนอตาราง ดังแสดง

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

Frequencies

LEVAGE					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	less or eq 40	4	2.9	2.9	2.9
	41-50 yrs	19	13.6	13.6	16.4
	51-60 yrs	26	18.6	18.6	35.0
	61-70 yrs	31	22.1	22.1	57.1
	Over 70	60	42.9	42.9	100.0
	Total	140	100.0	100.0	

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม EPI Info)

levage	Frequency	Percent	Cum Percent	
1	4	2.9%	2.9%	
2	19	13.6%	16.4%	
3	26	18.6%	35.0%	
4	31	22.1%	57.1%	
5	60	42.9%	100.0%	
Total	140	100.0%	100.0%	

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม MINITAB)

Tally for Discrete Variables: LEVAGE

LEVAGE	Count	Percent
1	4	2.86
2	19	13.57
3	26	18.57
4	31	22.14
5	60	42.86
N=	140	

ตารางที่ ... แสดงระดับอายุของผู้ป่วย จำนวน 140 ราย

ระดับอายุ	จำนวน (ราย)	ร้อยละ
ไม่เกิน 40 ปี	4	2.9
41-50 ปี	19	13.6
51-60 ปี	26	18.6
61-70 ปี	31	22.1
มากกว่า 70 ปี	60	42.9
รวม	140	100.0

ตัวอย่างการนำเสนอตารางแจกแจงความถี่แบบหลายทาง จำแนกตามระดับอายุของผู้ป่วย และประเภทการได้รับบริการพยาบาล (กลุ่มควบคุม และกลุ่มทดลอง) ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป และการนำเสนอตาราง ดังแสดง

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

Crosstabs

LEVAGE * Group of Intervention Crosstabulation

			Group of Intervention		Total
			Control (CRai)	Intervention (CM)	
LEVAGE	less or eq 40	Count	4	0	4
		% within Group of Intervention	5.7%	.0%	2.9%
	41-50 yrs	Count	7	12	19
		% within Group of Intervention	10.0%	17.1%	13.6%
	51-60 yrs	Count	17	9	26
		% within Group of Intervention	24.3%	12.9%	18.6%
	61-70 yrs	Count	15	16	31
		% within Group of Intervention	21.4%	22.9%	22.1%
	Over 70	Count	27	33	60
		% within Group of Intervention	38.6%	47.1%	42.9%
Total		Count	70	70	140
		% within Group of Intervention	100.0%	100.0%	100.0%

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม EPI Info)

GROUP

levage	0	1	TOTAL
1	4	0	4
Row %	100.0	0.0	100.0
Col %	5.7	0.0	2.9
2	7	12	19
Row %	36.8	63.2	100.0
Col %	10.0	17.1	13.6
3	17	9	26
Row %	65.4	34.6	100.0
Col %	24.3	12.9	18.6
4	15	16	31
Row %	48.4	51.6	100.0
Col %	21.4	22.9	22.1
5	27	33	60
Row %	45.0	55.0	100.0
Col %	38.6	47.1	42.9
TOTAL	70	70	140
Row %	50.0	50.0	100.0
Col %	100.0	100.0	100.0

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม MINITAB)

Tabulated Statistics: LEVAGE, GROUP

Rows: LEVAGE Columns: GROUP

	0	1	All
1	4 5.71	0 --	4 2.86
2	7 10.00	12 17.14	19 13.57
3	17 24.29	9 12.86	26 18.57
4	15 21.43	16 22.86	31 22.14
5	27 38.57	33 47.14	60 42.86
All	70 100.00	70 100.00	140 100.00

Cell Contents --

Count
% of Col

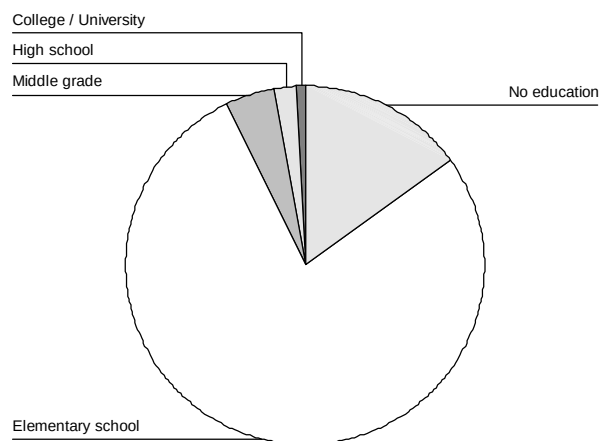
ตารางที่ ... แสดงระดับอายุของผู้ป่วย จำนวน 140 ราย

ระดับอายุ	ประเภทการได้รับการบริการ (ร้อยละ)	
	กลุ่มควบคุม	กลุ่มทดลอง
ไม่เกิน 40 ปี	4 (5.7)	-
41-50 ปี	7 (10.0)	12 (17.1)
51-60 ปี	17 (24.3)	9 (12.9)
61-70 ปี	15 (21.4)	16 (22.9)
มากกว่า 70 ปี	27 (38.6)	33 (47.1)
รวม	70 (100.0)	70 (100.0)

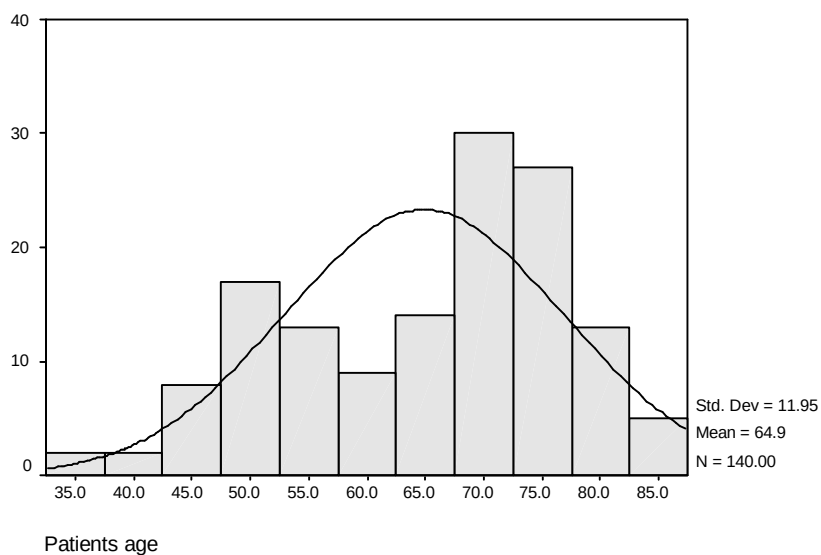
แผนภูมิ กราฟ

แบ่งเป็นการนำเสนอข้อมูลเชิงคุณภาพในรูปแบบแผนภูมิ กราฟ โดยทั่วไปนิยมใช้ กราฟวงกลม กราฟแท่ง ส่วนการนำเสนอข้อมูลเชิงปริมาณ นิยมนำเสนอเป็น ฮิสโตแกรม รูปหลายเหลี่ยมความถี่ กราฟแสดงความถี่สัมพัทธ์

ตัวอย่างการสร้างแผนภูมิ กราฟ โดยใช้โปรแกรม SPSS

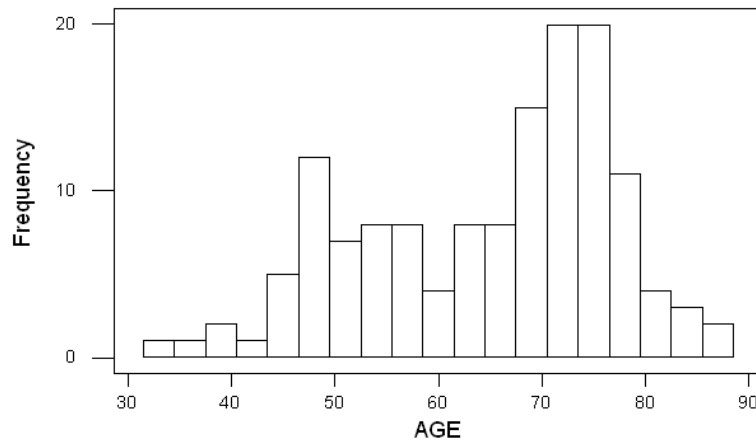


แผนภูมิวงกลม แสดงจำนวนและร้อยละของกลุ่มตัวอย่าง จำแนกตามระดับการศึกษา



ฮิสโตแกรม แสดงระดับอายุของผู้ป่วย จำนวน 140 ราย

ตัวอย่างการสร้างแผนภูมิ กราฟ โดยใช้โปรแกรม MINITAB



ฮิสโตแกรม แสดงระดับอายุของผู้ป่วย จำนวน 140 ราย

การนำเสนอตัวแทนของข้อมูล

การวัดแนวโน้มเข้าสู่ส่วนกลาง

ในการนำเสนอข้อมูลทั้งหมดข้างต้น ถ้ามีข้อมูลจำนวนมาก อาจทำให้ผู้อ่านหรือผู้ใช้เข้าใจได้ยาก จึงมีการนำเสนอค่าตัวแทนของข้อมูล เพื่อให้เห็นลักษณะของข้อมูลจากค่าตัวแทนของข้อมูลค่าใดค่าหนึ่ง ค่าตัวแทนของข้อมูล หรือการวัดแนวโน้มเข้าสู่ส่วนกลาง ที่นิยมใช้ได้แก่ ค่าเฉลี่ยเลขคณิต (arithmetic mean, $\bar{X}$) ค่ามัธยฐาน (median) และค่าฐานนิยม (mode)

ค่าเฉลี่ยเลขคณิต (arithmetic mean, $\bar{X}$)

$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n}$$

กรณีข้อมูลไม่ได้แจกแจงความถี่

$$\bar{X} = \frac{\sum_{i=1}^k f_i X_i}{n}$$

กรณีข้อมูลมีการแจกแจงความถี่

* กรณีข้อมูลมีการแจกแจงความถี่แบบจัดกลุ่ม (อันตรภาคชั้น) ให้ใช้จุดกึ่งกลาง (mid point) เป็นค่า X_i

ตัวอย่างการคำนวณค่าเฉลี่ยเลขคณิต

ตัวอย่างข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

10.6	9.9	12.6	15.2	12.3	11.7	12.3	12.5	11.8	12.4
10.2	11.3	9.4	11.4	11.0	11.6	12.2	13.4	9.9	9.8
10.2	9.2	15.3	10.9	9.0	11.0	8.6	12.5	11.6	12.6
16.7	7.7	10.9	10.1	8.7					

คำนวณจากข้อมูลจริง

$$\begin{aligned}\bar{X} &= (10.6 + 9.9 + 12.6 + 15.2 + \dots + 10.1 + 8.7) / 35 \\ &= 11.33\end{aligned}$$

คำนวณจากตารางแจกแจงความถี่

ตารางที่ ... แสดงระดับ RBC Cholinesterase (μ mole/นาที/ml) ในผู้ป่วย 35 ราย

ระดับ RBC Cholinesterase (μ mole/นาที/ml)	จุดกึ่งกลาง	จำนวน (ราย)	ร้อยละ
5.95 – 7.94	6.95	1	2.9
7.95 – 9.94	8.95	8	22.9
9.95 – 11.94	10.95	14	40.0
11.95 – 13.94	12.95	9	25.7
13.95 – 15.94	14.95	2	5.7
15.95 – 17.94	16.95	1	2.9
รวม		35	100.0

คำนวณจากตารางแจกแจงความถี่

$$\begin{aligned}\bar{X} &= [(6.95 \times 1) + (8.95 \times 8) + (10.95 \times 14) + \dots + (16.95 \times 1)] / 35 \\ &= 395.25 / 35 = 11.29\end{aligned}$$

คุณสมบัติของค่าเฉลี่ย

1. เป็นตัวแทนข้อมูล ที่ใช้ข้อมูลทุกค่ามาทำการคำนวณหาขนาดของค่าเฉลี่ย
2. เนื่องจากมีการนำข้อมูลทุกค่ามาคำนวณตามหลักคณิตศาสตร์ จึงสามารถใช้ในการวิเคราะห์สถิติขั้นสูงได้
3. เนื่องจากมีการใช้ข้อมูลทุกค่ามาคำนวณ ดังนั้นหากมีข้อมูลบางตัวที่มีขนาดใหญ่มากๆ หรือเล็กมากๆ ผิดปกติ จะมีผลต่อการคำนวณขนาดของค่าเฉลี่ยด้วย
4. ข้อมูลที่มีมาตรวัดเป็นนามบัญญัติ (nominal scale) และเรียงอันดับ (ordinal scale) ไม่สามารถใช้คำนวณค่าเฉลี่ยได้

ค่ามัธยฐาน (median)

เป็นค่าที่อยู่ตำแหน่งตรงกลางของข้อมูล เมื่อเรียงลำดับตามปริมาณข้อมูลทั้งหมด จากน้อยไปมาก หรือจากมากไปน้อย

ก. กรณีที่ข้อมูลไม่ได้แจกแจงความถี่

- ข้อมูลเป็นเลขคี่

$$\text{มัธยฐาน} = \text{ค่าของข้อมูลลำดับที่ } (n+1)/2$$

- ข้อมูลเป็นเลขคู่

$$\text{มัธยฐาน} = [\text{ค่าของข้อมูลลำดับที่ } (n)/2 + \text{ค่าของข้อมูลลำดับที่ } (n+1)/2] / 2$$

ข. กรณีที่ข้อมูลมีการแจกแจงความถี่

$$Me = L + I \left[\frac{n/2 - \sum f_L}{f_m} \right] \quad \text{หรือ} \quad Me = U + I \left[\frac{n/2 - \sum f_U}{f_m} \right]$$

โดยที่ L แทนขอบเขตล่างของชั้นมัธยฐาน

U แทนขอบเขตบนของชั้นมัธยฐาน

I แทนความกว้างของชั้น

n แทนจำนวนข้อมูลทั้งหมด

f_L แทนความถี่ของชั้นที่มีค่าสังเกตต่ำกว่าชั้นมัธยฐาน

f_U แทนความถี่ของชั้นที่มีค่าสังเกตสูงกว่าชั้นมัธยฐาน

f_m แทนความถี่ของชั้นที่มีค่ามัธยฐาน

ในที่นี้จะไม่กล่าวถึงการคำนวณในกรณีที่ข้อมูลมีการแจกแจงความถี่

ตัวอย่างการคำนวณค่ามัธยฐาน

จากข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

เมื่อนำข้อมูลมาเรียงลำดับจากน้อยไปมาก จะได้เป็น

7.7	8.6	8.7	9.0	9.2	9.4	9.8	9.9	9.9	10.1
10.2	10.2	10.6	10.9	10.9	11.0	11.0	11.3	11.4	11.6
11.6	11.7	11.8	12.2	12.3	12.3	12.4	12.5	12.5	12.6
12.6	13.4	15.2	15.3	16.7					

ข้อมูลที่อยู่ตรงกลางคือ ลำดับที่ $(n+1)/2 = (35+1)/2$ คือลำดับที่ 18

ค่ามัธยฐานของข้อมูลชุดนี้ คือ 11.3

คุณสมบัติของค่ามัธยฐาน

1. มัธยฐาน เป็นการใช้ค่าของข้อมูลที่อยู่ตำแหน่งตรงกลาง มาเป็นตัวแทน ดังนั้นข้อมูลที่มีค่ามาก หรือน้อยผิดปกติ จะไม่มีผลกระทบต่อค่ามัธยฐาน และถ้ามีการเปลี่ยนแปลงข้อมูลบางตัวในกลุ่มจะมีผลกระทบต่อค่ามัธยฐานน้อยมาก
2. มัธยฐาน จะเป็นค่าตัวแทนของข้อมูลได้ใกล้เคียงกับประชากรส่วนใหญ่มากกว่าค่าเฉลี่ย หากการแจกแจงข้อมูลเบ้ไปทางใดทางหนึ่ง
3. ข้อมูลที่มีมาตรวัดเป็นนามบัญญัติ (nominal scale) ไม่สามารถใช้คำนวณหาค่า มัธยฐานได้
4. กรณีที่มีข้อมูลกระจุกอยู่ที่ค่าต่ำสุด หรือสูงสุดมากเกินไป จะไม่สามารถหาค่ามัธยฐานได้ เช่น น้ำหนักของผู้ป่วย 9 คน เป็น 55 55 55 55 55 60 65 70 72 ค่ามัธยฐานเป็น 55 ซึ่งไม่ได้เป็นค่าของข้อมูลที่อยู่ครึ่งหนึ่งตามความหมายของมัธยฐาน

ฐานนิยม (mode)

เป็นค่าที่มีความถี่สูงสุดในข้อมูลชุดหนึ่ง ฐานนิยมอาจมีค่าเดียวในชุดข้อมูลนั้น หรืออาจมีหลายค่าได้ กรณีที่มีข้อมูลที่มีความถี่สูงสุดเท่ากันหลายค่า

กรณีที่มีข้อมูลมีการแจกแจงความถี่

$$Mo = L + I \left[\frac{d_1}{d_1 + d_2} \right]$$

โดยที่ L แทนขอบเขตล่างของชั้นที่มีความถี่สูงสุด

I แทนความกว้างของชั้น

d_1 แทนผลต่างระหว่างความถี่ของชั้นที่มีความถี่สูงสุด กับชั้นติดกันที่มีข้อมูลต่ำกว่า

d_2 แทนผลต่างระหว่างความถี่ของชั้นที่มีความถี่สูงสุด กับชั้นติดกันที่มีข้อมูลสูงกว่า

ตัวอย่างการคำนวณค่าฐานนิยม

จากข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

หากพิจารณาจากข้อมูลจริง จะพบว่าข้อมูลที่มีความถี่สูงสุดมีอยู่มากกว่า 1 รายการคือ คือ 9.9, 10.2, 10.9, 11.0, 11.6, 12.3, 12.5 และ 12.6

หากคำนวณจากตารางแจกแจงความถี่ จะได้ค่าฐานนิยม

$$Mo = L + I \left[\frac{d_1}{d_1 + d_2} \right] = 9.945 + 2 \left[\frac{6}{6 + 5} \right] = 11.036$$

คุณสมบัติของฐานนิยม

1. สามารถคำนวณได้ง่าย รวดเร็ว
2. ใช้กับข้อมูลที่มีมาตรวัดนามบัญญัติ (nominal scale)
3. ข้อมูลที่มีค่ามาก หรือน้อยผิดปกติ จะไม่มีผลกระทบต่อค่าฐานนิยม และถ้ามีการเปลี่ยนแปลงข้อมูลบางตัวในกลุ่มจะไม่มีผลกระทบต่อค่าฐานนิยม หรือมีน้อยมาก

ในการเลือกใช้การวัดแนวโน้มเข้าสู่ส่วนกลาง ตัวใด (ค่าเฉลี่ย มัธยฐาน ฐานนิยม) จะขึ้นกับ ลักษณะการกระจายของข้อมูล จุดประสงค์ของการนำไปใช้ และมาตรวัดของข้อมูลนั้นๆ

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

จากข้อมูล RBC Cholinesterase (m mole/นาท./ml) ในผู้ป่วย 35 ราย

Statistics			
RBC Cholinesterase			
N	Valid	35	จำนวนข้อมูลที่สมบูรณ์
	Missing	0	จำนวน Missing case
Mean		11.329	ค่าเฉลี่ย
Median		11.300	มัธยฐาน
Mode		9.9a	ฐานนิยม

a. Multiple modes exist. The smallest value is shown

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม EPI Info)

Obs	Total	Mean	Variance	Std Dev
35	396.5000	11.3286	3.7345	1.9325
Minimum	25%	Median	75%	Maximum
7.7000	9.9000	11.3000	12.4000	16.7000
				9.9000

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม MINITAB)

Results for: rbc.xls

Descriptive Statistics: RBC

Variable	N	Mean	Median	TrMean	StDev	SE Mean
RBC	35	11.329	11.300	11.232	1.932	0.327
Variable	Minimum	Maximum	Q1	Q3		
RBC	7.700	16.700	9.900	12.400		

การนำเสนอตำแหน่งหรือลำดับของข้อมูล

นอกจากการศึกษาการวัดแนวโน้มเข้าสู่ส่วนกลาง ซึ่งเป็นคุณลักษณะอย่างหนึ่งของข้อมูลแล้ว ยังอาจศึกษาลักษณะเฉพาะเจาะจงของค่าของข้อมูลบางค่าได้ เช่น ผู้ป่วยรายหนึ่งในกลุ่มตัวอย่าง มีรายได้ 72,000 บาทต่อปี อยากทราบว่าสถานะของผู้ป่วยคนนี้เป็นอย่างไร เมื่อพิจารณาเทียบรายได้ของผู้ป่วยที่เป็นกลุ่มตัวอย่างทั้งหมด ควรอยู่ในระดับสูงหรือต่ำอย่างไร หรืออยากทราบว่าผู้ป่วยในกลุ่มตัวอย่างเท่าไรที่มีรายได้สูงกว่า หรือต่ำกว่าผู้ป่วยรายนี้

การศึกษาดังกล่าวนี้นี้ เราเรียกว่าการกำหนดตำแหน่งของข้อมูล โดยทั่วไปมี 3 แบบคือ การกำหนดตำแหน่งของข้อมูลเป็น ควอร์ไทล์ (quartile) เดไซล์ (decile) และเปอร์เซ็นต์ไทล์ (percentile)

ควอร์ไทล์ (quartile)

เป็นการแบ่งข้อมูลออกเป็น 4 ส่วนเท่าๆ กัน เมื่อเรียงข้อมูลตามค่าของข้อมูลแล้ว เรียกค่าของข้อมูลที่ตรงกับจุดแบ่งข้อมูลดังกล่าวว่า ควอร์ไทล์ที่ 1, 2 และ 3

กรณีที่ข้อมูลไม่ได้แจกแจงความถี่

ข้อมูลในตำแหน่งควอร์ไทล์ที่ r (Q_r) จะอยู่ในตำแหน่งที่ $r(n+1)/4$

กรณีที่ข้อมูลแจกแจงความถี่

ข้อมูลในตำแหน่งควอร์ไทล์ที่ r (Q_r) จะอยู่ในตำแหน่งที่ $m/4$

เดไซล์ (decile)

เป็นการแบ่งข้อมูลออกเป็น 10 ส่วนเท่าๆ กัน เมื่อเรียงข้อมูลตามค่าของข้อมูลแล้ว เรียกค่าของข้อมูลที่ตรงกับจุดแบ่งข้อมูลดังกล่าวว่า เดไซล์ที่ 1, 2, 3, ..., 9

กรณีที่ข้อมูลไม่ได้แจกแจงความถี่

ข้อมูลในตำแหน่งเดไซล์ที่ r (D_r) จะอยู่ในตำแหน่งที่ $r(n+1)/10$

กรณีที่ข้อมูลแจกแจงความถี่

ข้อมูลในตำแหน่งเดไซล์ที่ r (D_r) จะอยู่ในตำแหน่งที่ $m/10$

เปอร์เซ็นต์ไทล์ (percentile)

เป็นการแบ่งข้อมูลออกเป็น 100 ส่วนเท่าๆ กัน เมื่อเรียงข้อมูลตามค่าของข้อมูลแล้ว เรียกค่าของข้อมูลที่ตรงกับจุดแบ่งข้อมูลดังกล่าวว่า เปอร์เซ็นต์ไทล์ ที่ 1, 2, 3, ..., 99

กรณีที่ข้อมูลไม่ได้แจกแจงความถี่

ข้อมูลในตำแหน่งเปอร์เซ็นต์ไทล์ ที่ r (Q_r) จะอยู่ในตำแหน่งที่ $r(n+1)/100$

กรณีที่ข้อมูลแจกแจงความถี่

ข้อมูลในตำแหน่งเปอร์เซ็นต์ไทล์ ที่ r (Q_r) จะอยู่ในตำแหน่งที่ $rn/100$

การคำนวณค่าของข้อมูลตำแหน่งที่ r ของควอร์ไทล์ เดซิซัล และเปอร์เซ็นต์ไทล์ กรณีที่ข้อมูลแจกแจงความถี่ ให้ใช้สูตรการคำนวณดังนี้

$$Q_r = L + I \left[\frac{rn/4 - \sum f_L}{f} \right]$$

$$D_r = L + I \left[\frac{rn/10 - \sum f_L}{f} \right]$$

$$P_r = L + I \left[\frac{rn/100 - \sum f_L}{f} \right]$$

โดยที่ L แทนขอบเขตล่างของชั้นที่ข้อมูลอยู่

I แทนความกว้างของชั้น

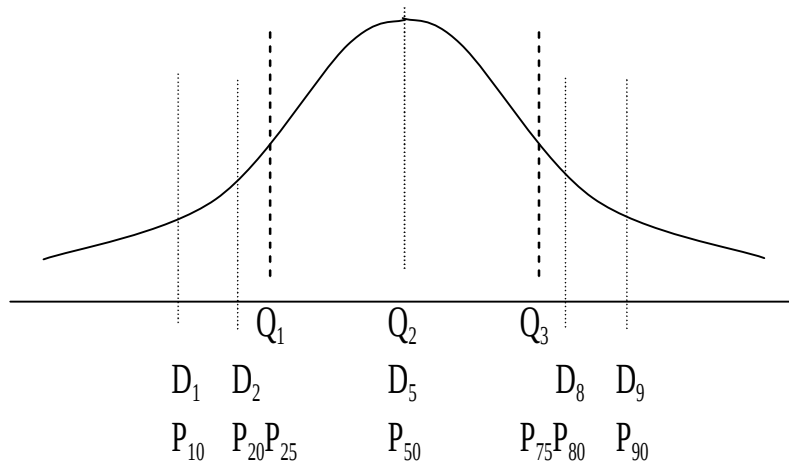
n แทนจำนวนข้อมูลทั้งหมด

f แทนความถี่ของชั้นที่ข้อมูลอยู่

f_L แทนความถี่ของชั้นที่มีค่าสังเกตต่ำกว่าชั้นที่ข้อมูลอยู่

ในที่นี้จะไม่กล่าวถึงการคำนวณในกรณีที่ข้อมูลมีการแจกแจงความถี่

ความสัมพันธ์ระหว่างตำแหน่งข้อมูลทั้ง 3 ชนิด



ตำแหน่งข้อมูลที่ Q1 จะตรงกับ P25

ตำแหน่งข้อมูลที่ Q2 จะตรงกับ D10, P50 และค่ามัธยฐาน

ส่วนตำแหน่งข้อมูลที่ Q3 จะตรงกับ P75

ตัวอย่างการคำนวณ

จากข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

เมื่อนำข้อมูลมาเรียงลำดับจากน้อยไปมาก จะได้เป็น

7.7	8.6	8.7	9.0	9.2	9.4	9.8	9.9	9.9	10.1
10.2	10.2	10.6	10.9	10.9	11.0	11.0	11.3	11.4	11.6
11.6	11.7	11.8	12.2	12.3	12.3	12.4	12.5	12.5	12.6
12.6	13.4	15.2	15.3	16.7					

- หา Q_1 เป็นข้อมูลตำแหน่งที่ $r(n+1)/4 = 1(35+1)/4 = 9$

ดังนั้นข้อมูลตำแหน่ง Q_1 คือ 9.9

- หา D_7 เป็นข้อมูลตำแหน่งที่ $r(n+1)/10 = 7(35+1)/10 = 25.2$

เนื่องจากข้อมูลตำแหน่งที่ 25 คือ 12.3 และตำแหน่งที่ 26 คือ 12.3 ดังนั้นข้อมูลตำแหน่ง D_7 คือ 12.3

- หา P_{90} เป็นข้อมูลตำแหน่งที่ $r(n+1)/100 = 90(35+1)/100 = 32.4$

เนื่องจากไม่มีข้อมูลตำแหน่งที่ 32.4 และข้อมูลตำแหน่งที่ 32 คือ 13.4 กับตำแหน่งที่ 33 คือ 15.2 มีค่าไม่เท่ากัน การหาค่าข้อมูลตำแหน่งที่ 32.4 จะต้องใช้การเทียบบัญญัติไตรยางค์ ดังนี้

ตำแหน่งต่างกัน $33-32 = 1$ ค่าของข้อมูลต่างกัน $15.2-13.4 = 1.8$

ตำแหน่งต่างกัน 0.4 ค่าของข้อมูลต่างกัน $(1.8 \times 0.4) / 1 = 0.72$

ดังนั้นข้อมูลตำแหน่ง P_{90} เท่ากับ $13.4 + 0.72 = 14.12$

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

จากข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

Statistics		
RBC Cholinesterase		
N	Valid	35
	Missing	0
Percentiles	10	8.880
	20	9.820
	25	9.900
	30	10.180
	40	10.900
	50	11.300
	60	11.660
	70	12.300
	75	12.400
	80	12.500
	90	14.120

การอธิบายการกระจายของข้อมูล

การกระจายของข้อมูลเป็นคุณลักษณะอย่างหนึ่งที่ใช้วัดความแตกต่างของค่าของข้อมูลทั้งหมด ข้อมูลที่มีการกระจายน้อย แสดงถึงข้อมูลมีการเกาะกลุ่มอยู่ที่ค่าใกล้เคียงกัน ส่วนข้อมูลที่มีการกระจายมาก แสดงว่าข้อมูลเกาะกลุ่มไม่ดี มีค่าแตกต่างกันมาก

ในการศึกษา เพื่อให้มองเห็นภาพของข้อมูลได้ชัดเจน จึงควรศึกษาการวัดแนวโน้มเข้าสู่ส่วนกลางควบคู่ไปกับการวัดการกระจายข้อมูล

วิธีวัดการกระจาย มีทั้งหมด 4 วิธี ได้แก่ พิสัย ส่วนเบี่ยงเบนควอร์ไทล์ ส่วนเบี่ยงเบนเฉลี่ย และส่วนเบี่ยงเบนมาตรฐาน

พิสัย (range)

เป็นความแตกต่างระหว่างข้อมูลที่มีค่าสูงสุด กับข้อมูลที่มีค่าต่ำสุด

พิสัย = ค่าสูงสุด - ค่าต่ำสุด

ส่วนเบี่ยงเบนควอร์ไทล์ (quartile deviation : QD)

หมายถึง ครึ่งหนึ่งของผลต่างระหว่างข้อมูลที่ Q_3 กับ Q_1

$$QD = [Q_3 - Q_1]/2$$

ส่วนเบี่ยงเบนเฉลี่ย (mean deviation : MD)

เป็นค่าเฉลี่ยความแตกต่างระหว่างข้อมูลทุกตัวกับค่าเฉลี่ย

$$MD = \frac{\sum_{i=1}^k f_i |x_i - \bar{X}|}{n}$$

ส่วนเบี่ยงเบนมาตรฐาน (standard deviation : S)

เป็นรากที่ 2 ของค่าเฉลี่ยของกำลังสองของผลต่างระหว่างข้อมูลทุกตัวกับค่าเฉลี่ย

- ส่วนเบี่ยงเบนมาตรฐานของประชากร

$$s = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{X})^2}{n}} \quad \text{หรือ} \quad s = \sqrt{\frac{\sum_{i=1}^n x_i^2}{n} - \bar{X}^2}$$

- ส่วนเบี่ยงเบนมาตรฐานของกลุ่มตัวอย่าง

$$S = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{X})^2}{n-1}} \quad \text{หรือ} \quad S = \sqrt{\frac{\sum_{i=1}^n x_i^2 - n\bar{X}^2}{n-1}}$$

และเรียกกำลังสองของส่วนเบี่ยงเบนมาตรฐาน ว่า ความแปรปรวน (variance) สัญลักษณ์เป็น s^2 หรือ S^2

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

จากข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

Statistics		
RBC		
N	Valid	35
	Missing	0
Std. Deviation		1.9325
Variance		3.7345
Range		9.00
Minimum		7.70
Maximum		16.70

จากผลการคำนวณ จะได้

ส่วนเบี่ยงเบนมาตรฐาน	1.9325	ความแปรปรวน	3.7345
พิสัย	9.00		
ค่าสูงสุด	7.70	ค่าต่ำสุด	16.70

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม MINITAB)

Descriptive Statistics: RBC

Variable	N	Mean	Median	TrMean	StDev	SE
Mean RBC 0.327	35	11.329	11.300	11.232	1.932	
Variable	Minimum	Maximum	Q1	Q3		
RBC	7.700	16.700	9.900	12.400		

การวัดการกระจายสัมพัทธ์

ใช้ในการเปรียบเทียบการกระจายของข้อมูลหลายชุด เนื่องจากข้อมูลที่ต่างชุดกันอาจมีธรรมชาติที่ต่างกัน เช่น การเปรียบเทียบการกระจายของน้ำหนักทารกแรกเกิด กับ การกระจายของข้อมูลน้ำหนักนักเรียนระดับประถมต้น ในกลุ่มนักเรียนระดับประถมต้นหากมีความแตกต่างของข้อมูล 1 กิโลกรัม ก็ไม่ถือว่ามีความแตกต่างกันมาก แต่ในขณะที่น้ำหนักทารกแรกเกิด ค่า 1 กิโลกรัมถือว่ามีความแตกต่างกันอย่างมาก ดังนั้นในการเปรียบเทียบการกระจายจึงไม่สามารถใช้การวัดการกระจายมาเปรียบเทียบกันได้ จะต้องใช้การกระจายสัมพัทธ์ ซึ่งได้แก่ สัมประสิทธิ์พิสัย สัมประสิทธิ์ส่วนเบี่ยงเบนควอร์ไทล์ สัมประสิทธิ์ส่วนเบี่ยงเบนเฉลี่ย และ สัมประสิทธิ์ส่วนเบี่ยงเบนมาตรฐาน

สัมประสิทธิ์พิสัย (coefficient of range, Coef.R)

$$\text{สัมประสิทธิ์พิสัย} = [\text{ค่าสูงสุด} - \text{ค่าต่ำสุด}] / [\text{ค่าสูงสุด} + \text{ค่าต่ำสุด}]$$

สัมประสิทธิ์ส่วนเบี่ยงเบนควอร์ไทล์ (coefficient of quartile deviation , Coef.QD)

$$\text{Coef.QD} = [Q_3 - Q_1] / [Q_3 + Q_1]$$

สัมประสิทธิ์ส่วนเบี่ยงเบนเฉลี่ย (coefficient of mean deviation , Coef.MD)

$$\text{Coef.MD} = \text{MD} / \bar{X}$$

สัมประสิทธิ์ส่วนเบี่ยงเบนมาตรฐาน (coefficient of standard deviation , Coef.S)

$$\text{Coef.S} = S / \bar{X}$$

โดยทั่วไปนิยมเรียกว่า สัมประสิทธิ์การแปรผัน (coefficient of variation , CV)

ตัวอย่าง

ศึกษาการเจริญเติบโตของเด็กกลุ่มหนึ่งจำนวน 200 คน โดยวัดค่าส่วนสูงและน้ำหนัก ได้ผลดังนี้

ส่วนสูง ค่าเฉลี่ย 142.7 ซม.	ส่วนเบี่ยงเบนมาตรฐาน 15.2
น้ำหนัก ค่าเฉลี่ย 38.8 กก.	ส่วนเบี่ยงเบนมาตรฐาน 6.5
CV ของส่วนสูง = $15.2/142.7$	= 0.1065 หรือ 10.65%
CV ของน้ำหนัก = $6.5/38.8$	= 0.1657 หรือ 16.57%

เนื่องจาก CV ของน้ำหนักมีค่ามากกว่า CV ของส่วนสูง แสดงว่า น้ำหนักมีการกระจายมากกว่าส่วนสูง

ผลการประเมินความวิตกกังวลของนักศึกษาพยาบาลปีที่ 4 จำนวน 40 คน และนักศึกษาพยาบาลปีที่ 3 จำนวน 60 คนมีดังนี้

นักศึกษาพยาบาลปีที่ 4 คะแนนเฉลี่ย 140 SD = 12

นักศึกษาพยาบาลปีที่ 3 คะแนนเฉลี่ย 160 SD = 28

CV ของนักศึกษาพยาบาลปีที่ 4 = $12/140 = .0857$ หรือ 8.57%

CV ของนักศึกษาพยาบาลปีที่ 3 = $28/160 = .175$ หรือ 17.5%

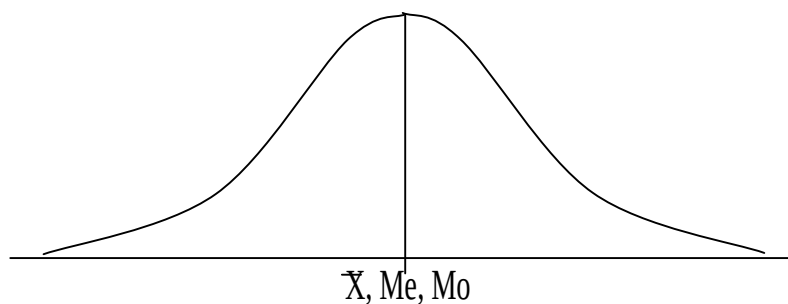
คะแนนการประเมินของนักศึกษาพยาบาลปีที่ 3 มีการกระจายเป็น $17.5/8.57 = 2.04$ เท่าของการประเมินของนักศึกษาพยาบาลปีที่ 4

การอธิบายการแจกแจงความถี่ของข้อมูล

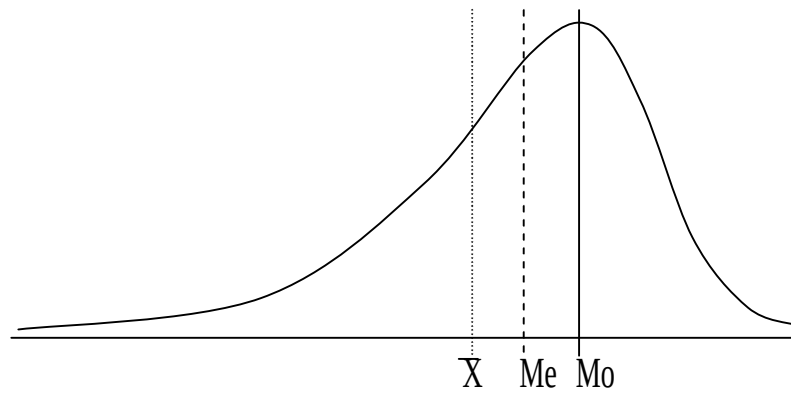
เป็นการศึกษาข้อมูลด้วยโค้งความถี่ของข้อมูล ค่าที่ใช้ในการวัดการแจกแจงข้อมูล แบ่งเป็น ความเบ้ของข้อมูล (skewness) และความโด่งของข้อมูล (kurtosis)

โค้งปกติ (normal curve)

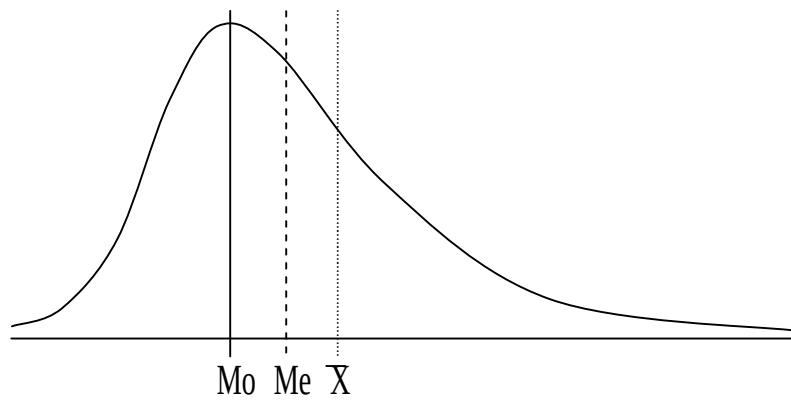
เป็นการแจกแจงความถี่ของข้อมูลต่อเนื่อง โดยเมื่อนำเสนอในรูปโค้งความถี่ จะได้โค้งสมมาตรมีลักษณะเป็นรูประฆังคว่ำ ดังรูป



ความเบ้ของข้อมูล

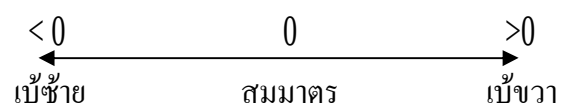


เบ้ซ้าย (เบ้ทางลบ) ค่าสัมประสิทธิ์ความเบ้เป็นค่าลบ



เบ้ขวา (เบ้ทางบวก) ค่าสัมประสิทธิ์ความเบ้เป็นค่าบวก

ค่าสัมประสิทธิ์ความเบ้



ความโด่งของข้อมูล

เป็นลักษณะความสูงของโค้งความถี่ จำแนกได้เป็น 3 แบบ คือ

โด่งปกติ (mesokurtic) เป็นลักษณะความสูงของโค้งสมมาตร ที่มีความสูงในระดับปานกลาง มีค่าสัมประสิทธิ์ความโด่ง = 0.263

โด่งสูง (leptokurtic) ความสูงของโค้งอยู่ในระดับสูง มีค่าสัมประสิทธิ์ความโด่ง > 0.263

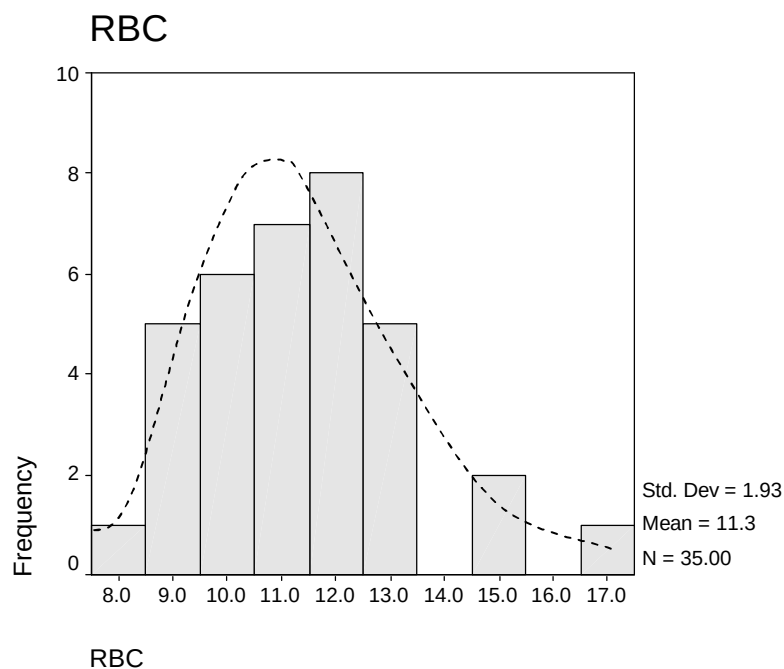
โด่งต่ำ (platykurtic) ความสูงของโค้งอยู่ในระดับต่ำ มีค่าสัมประสิทธิ์ความโด่ง < 0.263

ผลการวิเคราะห์โดยโปรแกรมสำเร็จรูป (โปรแกรม SPSS)

จากข้อมูล RBC Cholinesterase (m mole/นาที/ml) ในผู้ป่วย 35 ราย

Statistics

RBC		
N	Valid	35
	Missing	0
Skewness		.715
Std. Error of Skewness		.398
Kurtosis		.986
Std. Error of Kurtosis		.778



คะแนนมาตรฐาน

คะแนนมาตรฐาน (standard score) เป็นค่าคะแนนที่ได้จากการแปลงข้อมูลที่มีหน่วยวัด เช่น นาฬิกา เมตร กิโลกรัม ฯลฯ ให้เป็นค่าคะแนนที่มีคุณสมบัติคือไม่มีหน่วยวัด ทั้งนี้เพื่อให้สามารถนำข้อมูลที่มีหน่วยวัดที่แตกต่างกันมาเปรียบเทียบกันได้

คะแนนมาตรฐาน (standard score) หรือ Z-score
$$Z_i = \frac{x_i - \bar{X}}{S}$$

ตัวอย่าง

ศึกษาความวิตกกังวลของนักศึกษาระดับบัณฑิตศึกษา ในสาขาหนึ่ง จำนวน 50 คน โดยใช้แบบประเมินความวิตกกังวล แบบวัดดังกล่าวมี 2 ส่วน ส่วนที่ 1 คะแนนเต็ม 50 คะแนน ส่วนที่ 2 คะแนนเต็ม 80 คะแนน ผลการวัดคะแนนเป็นดังนี้

ส่วนที่ 1 คะแนนเต็ม 50 คะแนน ค่าเฉลี่ย 30 ส่วนเบี่ยงเบนมาตรฐาน 5

ส่วนที่ 2 คะแนนเต็ม 80 คะแนน ค่าเฉลี่ย 54 ส่วนเบี่ยงเบนมาตรฐาน 9

นาย ก. ได้คะแนนแบบวัดส่วนที่ 1 40 คะแนน และส่วนที่ 2 65 คะแนน อยากทราบว่า นาย ก. มีความวิตกกังวลในส่วนใดสูงกว่ากัน

จาก
$$Z_i = \frac{x_i - \bar{X}}{S}$$

ส่วนที่ 1 $Z = [40-30]/5 = 2.0$

ส่วนที่ 2 $Z = [65-54]/9 = 1.22$

แสดงว่า นาย ก. มีความวิตกกังวลในส่วนที่ 1 สูงกว่าส่วนที่ 2

สถิติอ้างอิง (inferential statistics)

จุดประสงค์หนึ่งที่สำคัญของการศึกษาสถิติ คือ การนำความรู้หรือข้อเท็จจริงที่ได้จากกลุ่มตัวอย่าง ซึ่งสุ่มมาจากประชากรที่ต้องการศึกษา ไปช่วยในการตัดสินใจหรือสรุปผลเกี่ยวกับคุณลักษณะของประชากร วิธีการดังกล่าวนี้เรียกว่า การอนุมาน (inferential) ทั้งนี้เนื่องจากการที่จะศึกษาคุณลักษณะของประชากรโดยใช้วิธีการรวบรวมข้อมูล ข้อเท็จจริงจากทั้งประชากรเป็นไปได้ยาก หรือไม่สามารถกระทำได้เลย

กลุ่มตัวอย่างที่ใช้ ในการอนุมานเชิงสถิติ จะต้องเป็นกลุ่มตัวอย่างที่ดี ซึ่งมีคุณสมบัติเป็นตัวแทนที่ดีของประชากร กระบวนการที่จะได้มาซึ่งตัวอย่างที่ดีนั้น เรียกว่า การสุ่มตัวอย่าง (sampling) ซึ่งจะมีวิธีการต่างๆ มากมาย ในที่นี้จะไม่ขอกล่าวถึง ค่าที่ได้จากกลุ่มตัวอย่าง เราเรียกว่า ค่าสถิติ ส่วนค่าที่ได้มาจากประชากร จะเรียกว่า พารามิเตอร์ ในการนำเอาค่าที่ได้จากกลุ่มตัวอย่าง (ค่าสถิติ) ไปสรุปผลเกี่ยวกับคุณลักษณะของประชากร (พารามิเตอร์) กระทำได้ 2 วิธีด้วยกัน คือ การประมาณค่า และการทดสอบสมมติฐาน

การประมาณค่า

หมายถึง วิธีการใช้ค่าสถิติที่ได้จากตัวอย่างไปประมาณค่าพารามิเตอร์ เป็นการหาข้อสรุปที่เกี่ยวกับพารามิเตอร์ ในลักษณะของการประมาณ ซึ่งมักแสดงในรูปตัวเลข เช่น ประมาณค่าเฉลี่ยของประชากร ประมาณค่าสัดส่วนของประชากร เป็นต้น

ในการประมาณค่าพารามิเตอร์ทางสถิติ มีวิธีการประมาณได้ 2 แบบ คือ การประมาณค่าแบบจุด (point estimation) และการประมาณค่าแบบช่วง (interval estimation)

ค่าที่ประมาณได้จากการประมาณค่าแบบจุด จะมีลักษณะเป็นตัวเลขประมาณค่าเดียว เรียกเป็น ค่าประมาณ (estimate) เช่น ประมาณค่าเฉลี่ยค่าใช้จ่ายในการรักษาพยาบาลของผู้ป่วยโรคเบาหวานที่มารับการรักษาในโรงพยาบาลมหาราชนครเชียงใหม่ เท่ากับ 4,950.-บาท

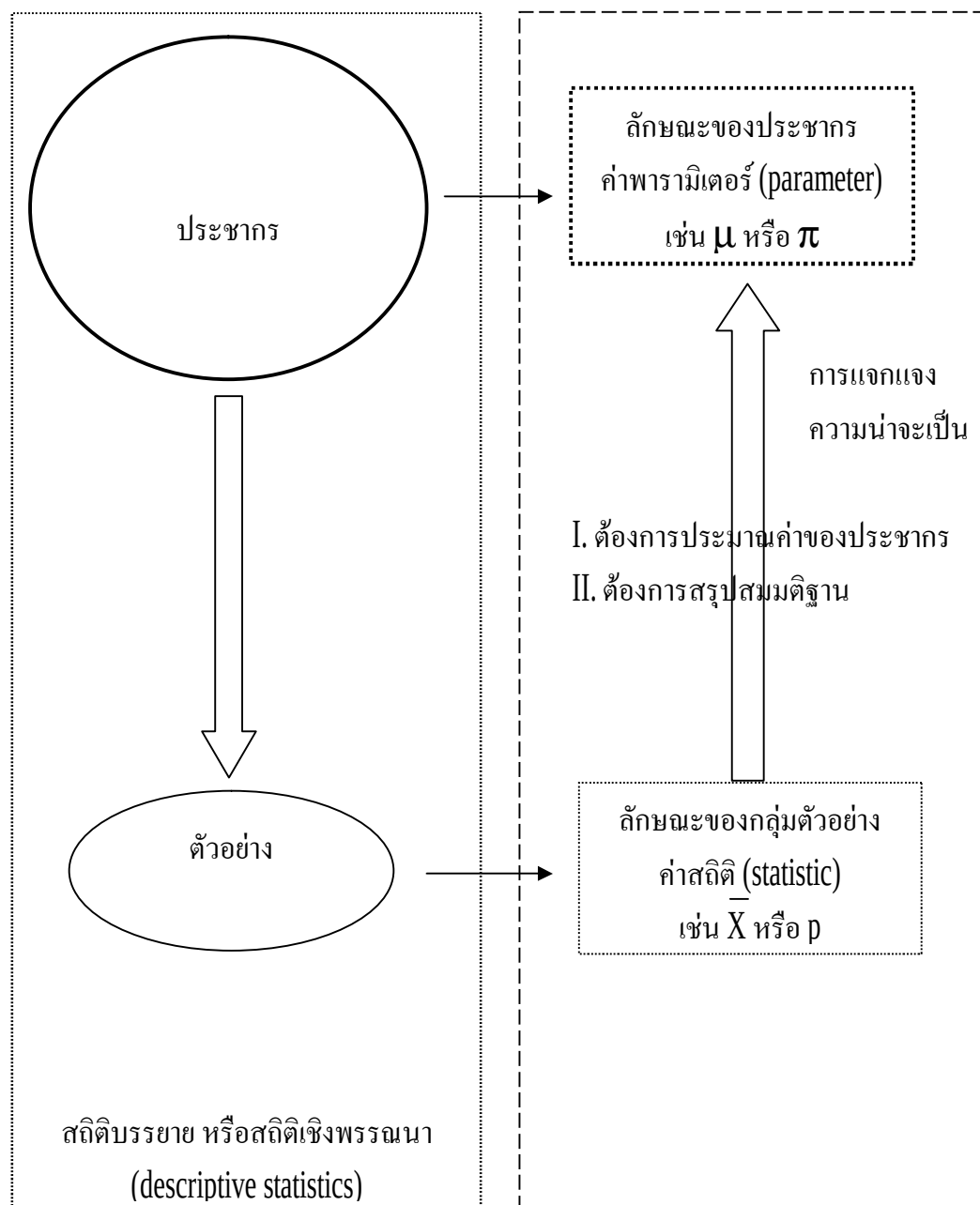
สำหรับค่าที่ประมาณได้จากการประมาณค่าแบบช่วง จะได้ช่วงของตัวเลขที่ประมาณ เรียก ช่วงการประมาณ เช่น ค่าเฉลี่ยค่าใช้จ่ายในการรักษาพยาบาลของผู้ป่วยโรคเบาหวานที่มารับการรักษาในโรงพยาบาลมหาราชนครเชียงใหม่ อยู่ระหว่าง 3,370 – 6,480.-บาท

การทดสอบสมมติฐาน

สมมติฐาน (hypothesis) คือ คำตอบที่คาดคะเนไว้ล่วงหน้า และคำตอบนี้ได้มาจากหลักการทางเหตุผล ซึ่งมาจากความรู้เดิม ประสบการณ์ เอกสาร ตำรา หรือทฤษฎีที่เกี่ยวข้อง

สมมติฐานการวิจัย คือ ความเชื่อของผู้วิจัยว่าเรื่องที่สนใจศึกษาจะมีลักษณะอย่างใดอย่างหนึ่ง ความเชื่อนั้นจะเป็นจริงหรือไม่ก็ได้ เช่น ผู้วิจัยเชื่อว่ายาแก้ปวด A สามารถลดความเจ็บปวดได้อย่างมีประสิทธิภาพ หรือ ผู้ป่วยผ่าตัดเปลี่ยนลิ้นหัวใจที่ได้รับคำแนะนำในการปฏิบัติตัวจะมีระดับความวิตกกังวลต่ำกว่าที่ไม่ได้รับคำแนะนำ

การทดสอบความเชื่อหรือสิ่งที่คาดไว้ เรียกว่า การทดสอบสมมติฐานทางสถิติ



จากที่กล่าวมาการใช้สถิติเชิงอนุมาน เป็นการประมาณค่า หรือสรุปผลการทดสอบ สมมติฐานจากข้อมูลที่ได้มาจากกลุ่มตัวอย่าง เพื่อไปสรุปหรือคาดการณ์ถึงลักษณะของประชากร ดังนั้นในกระบวนการดังกล่าวจึงเกี่ยวข้องกับ ความน่าจะเป็นของการเกิดเหตุการณ์ต่างๆ เราจึงต้องทราบถึงลักษณะของการแจกแจงความน่าจะเป็นที่เกี่ยวข้องกับการประมาณค่า และการทดสอบ สมมติฐานทางสถิติ ดังนี้

ความน่าจะเป็น และการแจกแจงความน่าจะเป็น

ความน่าจะเป็นของการเกิดเหตุการณ์ใด คือ โอกาสที่จะเกิดเหตุการณ์นั้น จากการเกิด เหตุการณ์ต่างๆ ทั้งหมด วัตถุประสงค์ออกมาในรูปของตัวเลข

ให้ E แทนเหตุการณ์ที่สนใจ

$n(E)$ แทนจำนวนครั้งของการเกิดเหตุการณ์ที่สนใจ

$n(S)$ แทนจำนวนครั้งของการเกิดเหตุการณ์ทั้งหมด

และ $P(E)$ แทนความน่าจะเป็นของการเกิดเหตุการณ์ที่สนใจ

จะได้ว่า $P(E) = n(E) / n(S)$

ตัวอย่างเช่น สํารวจประชาชนตำบลหนึ่งในจังหวัดเชียงใหม่ จำนวน 2500 คน พบว่า ป่วยเป็น ไข้เลือดออกจำนวน 25 ราย

ความน่าจะเป็นที่ประชาชนในตำบลนี้จะเป็นไข้เลือดออก = $25/2500 = 0.01$

การแจกแจงความน่าจะเป็น เป็นการหาความน่าจะเป็นที่เกิดขึ้นกับเหตุการณ์ทุกกรณีที่เป็นไปได้ โดยอาศัยความถี่ของข้อมูลจำนวนมาก

การแจกแจงความน่าจะเป็น ที่สำคัญๆ และมีเกี่ยวข้องกับการวิเคราะห์ข้อมูลทางสถิติ ได้แก่ การแจกแจงทวินาม การแจกแจงปัวซอง การแจกแจงแบบปกติ เป็นต้น

การแจกแจงทวินาม (Binomial distribution)

การแจกแจงทวินาม มาจากการทดลองที่เรียกว่า Bernoulli trial เป็นการทดลองที่สนใจ ผลลัพธ์เพียง 2 อย่าง คือ เกิดผลเป็นไปตามที่กำหนด เรียก success หรือไม่เกิดผลตามที่กำหนด เรียก failure

การแจกแจงทวินาม จะเป็นการแจกแจงความน่าจะเป็นที่เกิดจากการทดลองแบบ bernoulli ซ้ำกันหลายๆ ครั้ง

เช่น สุ่มตัวอย่างประชาชนในจังหวัดเชียงใหม่ 200 คน สนใจคนที่มึนเลือดหมู A จำนวน ไม่ต่ำกว่า 25 คน

การแจกแจงปัวซอง (Poisson distribution)

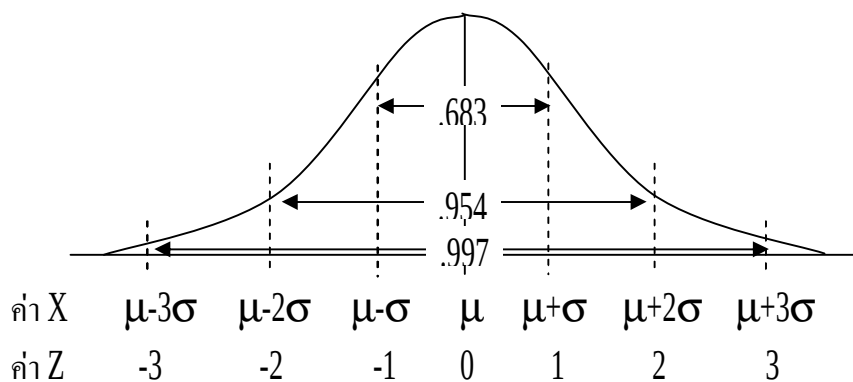
เป็นการแจกแจงความน่าจะเป็น ที่เกิดจากการทดลองปัวซอง เป็นการทดลองที่เกิดขึ้นในช่วงเวลาหนึ่ง หรือขอบเขตหนึ่ง โดยที่เหตุการณ์ที่สนใจมีได้ตั้งแต่ไม่เกิดขึ้นเลย จนถึงเกิดขึ้นอย่างไม่จำกัด

การแจกแจงแบบปกติ (Normal distribution)

เป็นการแจกแจงความน่าจะเป็นของข้อมูลต่อเนื่อง มีความสำคัญและมีการใช้มากที่สุด ในทางสถิติ ลักษณะเป็นโค้งระฆังคว่ำ สมมาตร

คุณสมบัติ

1. ลักษณะโค้งระฆังคว่ำ สมมาตร ปลายโค้งทั้งสองข้างจะค่อยๆ ลาดไปสู่แกนอน ที่ระยะอนันต์ แต่จะไม่ตัดแกน
2. พื้นที่ทั้งหมดระหว่างเส้นโค้งกับแกนอน จะเป็นตารางหน่วย พื้นที่ใต้โค้งปกติแทนความน่าจะเป็นของเหตุการณ์ทั้งหมด ซึ่งมีค่าเท่ากับ 1
3. เส้นโค้งจะมีจุดสูงสุดที่ค่าเฉลี่ย หากใช้ตำแหน่งค่าเฉลี่ยเป็นตัวแบ่งจะแบ่งข้อมูลเป็น 2 ด้าน ด้านขวาจะเป็นด้านที่มีค่าสูงกว่าค่าเฉลี่ย ส่วนด้านซ้ายจะเป็นด้านที่มีค่าต่ำกว่าค่าเฉลี่ย
4. ค่าส่วนเบี่ยงเบนมาตรฐานจะแบ่งพื้นที่ภายใต้โค้งปกติออกเป็น 6 ส่วน ดังรูป



ชนิดของการทดสอบสมมติฐานทางสถิติ

สมมติฐานทางสถิติ แบ่งเป็น สมมติฐานเพื่อการทดสอบ (null hypothesis) และ สมมติฐานแย้ง (alternative hypothesis)

สมมติฐานเพื่อการทดสอบ (null hypothesis) เขียนแทนด้วยสัญลักษณ์ H_0 เป็นสมมติฐานที่มีลักษณะเป็นการกำหนดค่าพารามิเตอร์ที่แน่นอน ต้องการทดสอบว่าเป็นความจริงหรือไม่

สมมติฐานแย้ง (alternative hypothesis) เขียนแทนด้วยสัญลักษณ์ H_1 เป็นสมมติฐานที่ตั้งขึ้นมาควบคู่กับ H_0 เพื่อเป็นทางเลือกหรือข้อแย้งกับ H_0 ในกรณีที่ต้องปฏิเสธ H_0

การตั้งสมมติฐานทางสถิติ จะตั้งสมมติฐานทั้ง H_0 และ H_1 ควบคู่กันเสมอ ตัวอย่างการตั้งสมมติฐาน เช่น

$$\begin{array}{ll} H_0 : \mu = 2100 & H_1 : \mu \neq 2100 \\ \text{หรือ} & H_0 : \mu_1 = \mu_2 & H_1 : \mu_1 > \mu_2 \end{array}$$

แนวความคิดในการทดสอบสมมติฐาน

การทดสอบสมมติฐานเป็นกระบวนการที่จะนำไปสู่การตัดสินใจหรือสรุปผล โดยอยู่บนพื้นฐานของหลักฐานที่ได้จากตัวอย่าง ดังนั้นในการตัดสินใจจึงอาจมีความผิดพลาดได้ ซึ่งความผิดพลาดที่เกิดจากการตัดสินใจ หรือสรุปผล เรียกว่า ความคลาดเคลื่อน (error) อาจเกิดขึ้นได้ 2 ชนิด คือ ความคลาดเคลื่อนชนิดที่ 1 (type I error) และความคลาดเคลื่อนชนิดที่ 2 (type II error)

ความคลาดเคลื่อนชนิดที่ 1 (type I error) เป็นความคลาดเคลื่อนที่เกิดจากการตัดสินใจที่ปฏิเสธสมมติฐาน H_0 ทั้งๆที่สมมติฐาน H_0 ถูกต้องเป็นจริง มักแทนความน่าจะเป็นที่จะเกิดความคลาดเคลื่อนชนิดที่ 1 ด้วยสัญลักษณ์ α

ความคลาดเคลื่อนชนิดที่ 2 (type II error) เป็นความคลาดเคลื่อนที่เกิดจากการตัดสินใจที่ยอมรับสมมติฐาน H_0 ทั้งๆที่สมมติฐาน H_0 ไม่เป็นจริง มักแทนความน่าจะเป็นที่จะเกิดความคลาดเคลื่อนชนิดที่ 2 ด้วยสัญลักษณ์ β

สรุปการตัดสินใจที่จะเกิดขึ้นได้ในการทดสอบสมมติฐาน ดังตาราง

		ความเป็นจริงของ H_0	
		H_0 เป็นจริง	H_0 เป็นเท็จ
การตัดสินใจ	ยอมรับ H_0	$\bar{u}$ (1- α)	Type II error β
	ปฏิเสธ H_0	Type I error α	$\bar{u}$ (1- β)

จะเรียก ความน่าจะเป็นในการตัดสินใจว่าสมมติฐาน H_0 เป็นเท็จ ว่า อำนาจการทดสอบ (power of testing) และเรียก ความน่าจะเป็นของการปฏิเสธ H_0 เมื่อ H_0 เป็นจริง ว่า ระดับนัยสำคัญ (level of significant)

ในการทดสอบสมมติฐานจะเกิดความคลาดเคลื่อนทั้งสองแบบ เพื่อให้การทดสอบสมมติฐานได้ผลที่มีความถูกต้องและน่าเชื่อถือ จะต้องทำให้เกิดความคลาดเคลื่อนทั้งสองแบบนี้มีค่าน้อยที่สุด กล่าวคือต้องให้ค่า α และ β มีค่าน้อยๆ อย่างไรก็ตามเราไม่สามารถทำให้ค่า α และ β มีค่าน้อยลง หรือเป็นศูนย์ไปพร้อมกันทั้งสองค่าได้ ถ้า α มีค่าน้อย จะทำให้ค่า β มีค่ามาก และถ้า α มีค่ามาก จะทำให้ค่า β มีค่าน้อย ดังนั้นในทางทฤษฎีจึงกำหนดค่า α ให้เป็นมาตรฐาน และพยายามหาวิธีการทำให้ค่า β มีค่าน้อยที่สุด โดยทั่วไปมี 2 ทางคือ การเพิ่มขนาดตัวอย่าง และการพิจารณาตัวสถิติที่ดีที่สุดในการทดสอบ

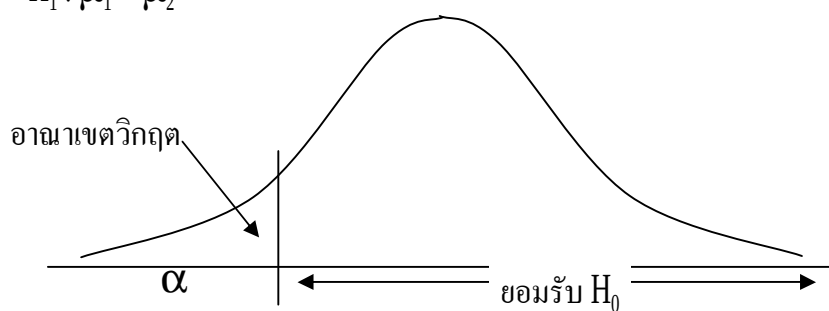
ประเภทของการทดสอบสมมติฐานทางสถิติ

ในการทดสอบสมมติฐาน จะกำหนดแบ่งพื้นที่การแจกแจงความน่าจะเป็นออกเป็น 2 ส่วนคือ อาณาเขตการยอมรับสมมติฐาน (acceptance region) และอาณาเขตปฏิเสธสมมติฐาน (rejection region) หรือเรียกว่า อาณาเขตวิกฤต (critical region) โดยพื้นที่ในอาณาเขตวิกฤต จะมีค่าเท่ากับความน่าจะเป็นของการปฏิเสธ H_0 เมื่อ H_0 เป็นจริง (คือระดับนัยสำคัญ : ค่า α) การกำหนดอาณาเขตวิกฤต ดังกล่าว หากมีการกำหนดพื้นที่ทั้งสองด้าน ของการแจกแจงความน่าจะเป็น จะเรียกว่า การทดสอบแบบสองทาง (two-tailed test) หากมีการกำหนดพื้นที่ด้านใดด้านหนึ่งของการแจกแจงความน่าจะเป็น จะเรียกว่า การทดสอบแบบทางเดียว (one-tailed test) ดังรูป

การทดสอบแบบทางเดียว (one-tailed test)

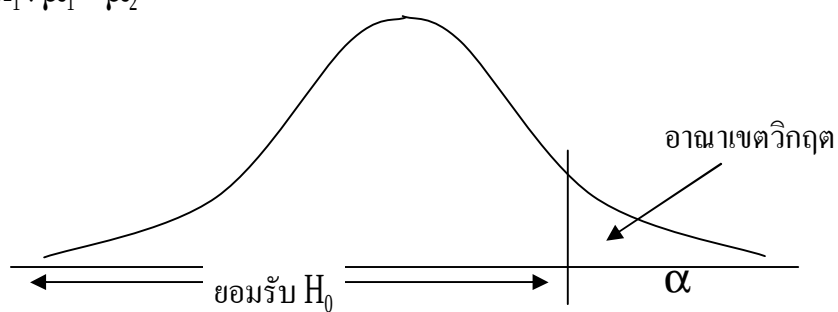
$$H_0 : \mu_1 = \mu_2$$

$$H_1 : \mu_1 < \mu_2$$



$$H_0 : \mu_1 = \mu_2$$

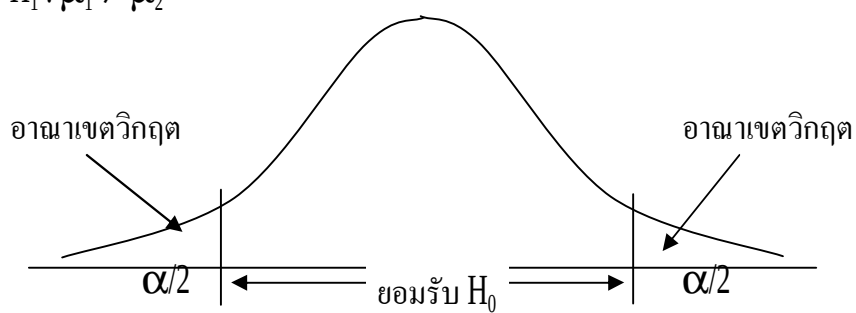
$$H_1 : \mu_1 > \mu_2$$



การทดสอบแบบสองทาง (two-tailed test)

$$H_0 : \mu_1 = \mu_2$$

$$H_1 : \mu_1 \neq \mu_2$$



ในการทดสอบสมมติฐานทางสถิติ ข้อมูลที่ได้จากการศึกษาในตัวอย่างจะนำมาใช้คำนวณค่าความน่าจะเป็นที่เรียกว่า p-value เพื่อใช้ตรวจสอบกับระดับนัยสำคัญทางสถิติ ในการสรุปการทดสอบสมมติฐาน

ถ้า $p\text{-value} \leq$ ระดับนัยสำคัญ (α) จะปฏิเสธ H_0 และเรียกว่า มีนัยสำคัญทางสถิติ

ถ้า $p\text{-value} >$ ระดับนัยสำคัญ (α) จะไม่ปฏิเสธ H_0 (ยอมรับ H_0) เรียกว่า ไม่มีนัยสำคัญทางสถิติ

กระบวนการดังกล่าว หากเป็นการคำนวณโดยผู้วิเคราะห์ทำการคำนวณเองโดยอาศัยนิยามและสูตรทางสถิติ จะมีข้อจำกัดในการหาพื้นที่ความน่าจะเป็น ดังนั้นในทางปฏิบัติจึงใช้ค่าสถิติที่คำนวณได้ ไปเปรียบเทียบกับค่าสถิติจากตารางสำเร็จรูป ณ ระดับนัยสำคัญที่กำหนด ดังนั้น แล้วจึงทำการสรุป ตามเกณฑ์ที่กำหนดของตัวสถิตินั้นๆ เช่น

ในการทดสอบค่าเฉลี่ย โดยใช้สถิติ Z

ถ้าค่า $|Z|$ ที่คำนวณ $\geq$ ค่า Z จากตาราง จะปฏิเสธ H_0 และเรียกว่า มีนัยสำคัญทางสถิติ

ถ้าค่า $|Z|$ ที่คำนวณ $<$ ค่า Z จากตาราง จะไม่ปฏิเสธ H_0 (ยอมรับ H_0) เรียกว่า

ไม่มีนัยสำคัญทางสถิติ

ขั้นตอนการทดสอบสมมติฐาน

ขั้นที่ 1 กำหนดสมมติฐานเพื่อการทดสอบ (H_0) และสมมติฐานแย้ง (H_1)

ขั้นที่ 2 พิจารณาเลือกตัวสถิติที่ใช้ในการทดสอบ ที่เหมาะสมกับพารามิเตอร์ที่ต้องการทดสอบ

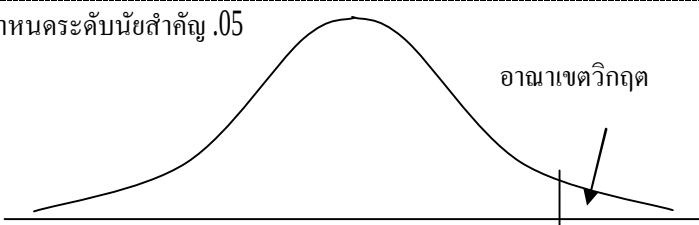
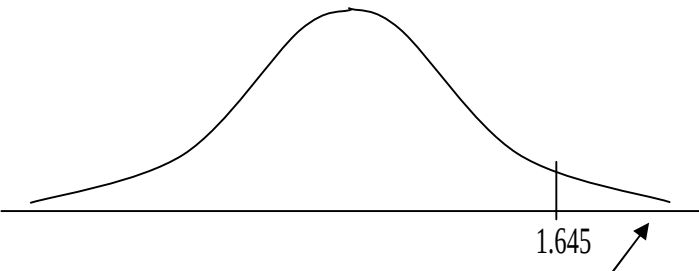
ขั้นที่ 3 กำหนดระดับนัยสำคัญ (α) และหาอาณาเขตวิกฤตของตัวสถิติที่ใช้ในการทดสอบภายใต้สมมติฐานที่กำหนด

ขั้นที่ 4 คำนวณค่าสถิติทดสอบจากข้อมูล

ขั้นที่ 5 เปรียบเทียบค่าสถิติทดสอบที่คำนวณได้กับค่าวิกฤต (ที่ได้จากการเปิดตารางค่าวิกฤตตามชนิดของสถิติต่าง) ถ้าค่าสถิติทดสอบตกอยู่ในอาณาเขตวิกฤต จะสรุปว่าปฏิเสธ H_0

ตัวอย่างขั้นตอนการทดสอบสมมติฐาน

จากการประเมินการปฏิบัติการดูแลตนเองของผู้ป่วยโรคความดันโลหิตสูง จำนวน 40 ข้อ ประเมินจากกลุ่มตัวอย่างผู้ป่วยโรคความดันโลหิตสูง 100 คน ได้คะแนนเฉลี่ย 90 คะแนน และจากการศึกษาที่ผ่านมาพบว่าส่วนเบี่ยงเบนมาตรฐานของประชากร เท่ากับ 5 คะแนน จะสรุปได้หรือไม่ว่าโดยส่วนใหญ่ผู้ป่วยมีการดูแลตนเองในระดับสูงกว่ามาตรฐาน (กำหนดให้มาตรฐานการดูแลตนเองจะต้องไม่ต่ำกว่า 80 คะแนน)

ขั้นที่ 1 กำหนดสมมติฐาน	$H_0 : \mu = 80$ $H_1 : \mu > 80$
ขั้นที่ 2 พิจารณาเลือกตัวสถิติที่ใช้ในการทดสอบ	พารามิเตอร์ของประชากร คือ ค่าเฉลี่ย จำนวนตัวอย่างมีจำนวนมาก $n = 100$ และทราบค่า σ ใช้สถิติทดสอบ Z
ขั้นที่ 3 กำหนดระดับนัยสำคัญ (α) และหาอาณาเขตวิกฤต	กำหนดระดับนัยสำคัญ .05 
ขั้นที่ 4 คำนวณค่าสถิติทดสอบจากข้อมูล	$n = 100, \bar{X} = 90, S = 5$ $Z = \frac{\bar{X} - m}{S / \sqrt{n}} = \frac{90 - 80}{5 / \sqrt{100}} = 20$
ขั้นที่ 5 เปรียบเทียบค่าสถิติทดสอบที่คำนวณได้กับค่าวิกฤต	จากตารางค่า Z ที่ $\alpha = .05$ มีค่า 1.645  ค่า Z จากการคำนวณ = 20 ตกในอาณาเขตวิกฤต ปฏิเสธ H_0 ค่าเฉลี่ยคะแนนการดูแลตนเองของผู้ป่วยกลุ่มนี้สูงกว่าระดับมาตรฐานอย่างมีนัยสำคัญทางสถิติ ที่ระดับ .05

แบบฝึกหัด

1. ในการนำเลือดของผู้ป่วยประมาณ 10 ลบ.ซม. มาตรวจเพื่อวิเคราะห์ว่าผู้ป่วยคนนั้นเป็นโรคเอดส์หรือไม่ อยากทราบว่าประชากร กลุ่มตัวอย่าง และพารามิเตอร์ของปัญหานี้คืออะไร
.....
.....
2. จงหาค่า Mean และ Median ของคะแนนต่อไปนี้
24 12 16 10 13 14 18 15 8 21 17
.....
.....
3. ในการทำวิจัยเรื่องหนึ่ง ผู้วิจัยต้องการเลือกกลุ่มตัวอย่างที่มีความวิตกกังวลใกล้เคียงกันมาทำการทดลอง ปรากฏว่ามีกลุ่มตัวอย่าง 2 กลุ่มให้เลือก คือ กลุ่ม ก และ กลุ่ม ข โดยทั้งสองกลุ่มมีค่าเฉลี่ยระดับความวิตกกังวลเท่ากัน คือ 86 คะแนน และกลุ่ม ก มีส่วนเบี่ยงเบนมาตรฐานเป็น 7 คะแนน และ กลุ่ม ข มีส่วนเบี่ยงเบนมาตรฐานเป็น 5 คะแนน ท่านจะเลือกกลุ่มใดเป็นกลุ่มตัวอย่างในการทำวิจัยเพราะเหตุใด
.....
.....
4. คะแนนจากการประเมินการปฏิบัติการดูแลตนเองของผู้ป่วยโรคความดันโลหิตสูง จำนวน 50 ข้อ ประเมินจากผู้ป่วยโรคความดันโลหิตสูง 50 คน ได้คะแนนเฉลี่ย 85 คะแนน ส่วนเบี่ยงเบนมาตรฐาน 5 คะแนน จงหาคะแนนมาตรฐานของผู้ป่วยที่ได้คะแนนการปฏิบัติในการดูแลตนเอง 80 คะแนน
.....
.....
.....
.....

5. การวัดระดับ Cholesterol ในเลือดของผู้สูงอายุ จำนวน 30 คน ปรากฏผล ดังนี้

260	179	198	297	305	250	270	187	172	300
197	146	138	268	288	277	195	289	193	205
215	224	268	282	235	249	222	256	273	284

- ก. จงสร้างตารางแจกแจงความถี่โดยให้อันตรภาคชั้นเป็น 10 พร้อมทั้งหาขอบเขตบน-ล่าง และจุดกลาง

- ข. จงคำนวณค่า Mean, Mode, Median

.....

.....

.....

.....

6. คะแนนจากการวัดระดับความพึงพอใจในงานของพยาบาล โรงพยาบาลแห่งหนึ่งปรากฏผล ดังนี้

หอผู้ป่วยศัลยกรรม	75	90	82	65	88	68	73	80	70	77
หอผู้ป่วยอายุรกรรม	83	76	72	87	92	85	83	79	90	90
หอผู้ป่วยสูติกรรม	69	84	86	85	79	80	79	84	78	83

จงคำนวณหา

- ก. Median ของระดับความพึงพอใจในงานของพยาบาล หอผู้ป่วยศัลยกรรม
- ข. Mean ของระดับความพึงพอใจในงานของพยาบาล หอผู้ป่วยอายุรกรรม
- ค. Mode ของระดับความพึงพอใจในงานของพยาบาล หอผู้ป่วยสูติกรรม
- ง. ความแปรปรวนและส่วนเบี่ยงเบนมาตรฐานของระดับความพึงพอใจของพยาบาล
ทุกชั้น
- จ. ข้อมูลใดมีการกระจายของข้อมูลมากที่สุด พิจารณาได้จากค่าใด

.....

.....

.....

.....

.....

7. น้ำหนักเฉลี่ยของเด็กหญิงกลุ่มหนึ่งเท่ากับ 22 กิโลกรัม ส่วนเบี่ยงเบนมาตรฐานเท่ากับ 5 กิโลกรัม น้ำหนักเฉลี่ยของเด็กชายกลุ่มหนึ่งเท่ากับ 24 กิโลกรัม ส่วนเบี่ยงเบนมาตรฐานเท่ากับ 3 กิโลกรัม ลักษณะการกระจายของน้ำหนักของเด็กหญิงหรือเด็กชายมากกว่ากัน พิจารณาจากค่าใด

.....

.....

8. จากการสำรวจอายุผู้ป่วยด้วยโรคมะเร็ง ณ โรงพยาบาลแห่งหนึ่งผลการสำรวจดังตารางต่อไปนี้

อายุผู้ป่วย (ปี)	17 – 26	27 – 36	37 – 46	47 – 56	57 – 66	67 เป็นต้นไป
จำนวนผู้ป่วย	3	4	14	18	6	5

ก. จงวัดการกระจายของอายุผู้ป่วยด้วยโรคมะเร็ง

ข. จงหาค่ากลางที่เหมาะสมของผู้ป่วยด้วยโรคมะเร็ง

.....

.....

.....

9. จากข้อมูลซึ่งรวบรวมมาจากการวัดความดันโลหิตของผู้ป่วยจำนวน 30 คน ได้ดังนี้

120 115 108 111 116 132 98 107 114 118 126 124 117 102
 105 123 135 120 128 129 127 130 114 126 133 120 128 135
 110 96

ก. จงสร้างตารางแจกแจงความถี่โดยให้ความกว้างของอันตรภาคชั้นเท่ากับ 5 โดยให้เริ่มจากค่าความดันที่วัดได้ต่ำสุด

ข. จงเขียนกราฟฮิสโทแกรม

ค. จงเปรียบเทียบค่าเฉลี่ยเลขคณิตที่คำนวณได้จากข้อมูลดิบและจากตารางแจกแจงความถี่

ง. จงเปรียบเทียบค่ามัธยฐานที่คำนวณได้จากข้อมูลดิบและจากตารางแจกแจงความถี่

.....

.....

.....

10. ในการบันทึกความยาวและน้ำหนักของทารกแรกเกิด จำนวน 45 คน ปรากฏผลดังตาราง

ความยาว (ซม.)	จำนวนทารก	น้ำหนัก (กรัม)	จำนวนทารก
ไม่เกิน 30	3	1501 – 2000	4
31 – 35	6	2001 – 2500	9
36 – 40	7	2501 – 3000	10
41 – 45	8	3001 – 3500	13
46 – 50	19	3501 – 4000	4
ตั้งแต่ 51 ขึ้นไป	2	ตั้งแต่ 4001 เป็นต้นไป	5
	45		45

ก. จงใช้เทคนิคทางสถิติที่เหมาะสมวัดค่ากลางของความยาวทารกแรกเกิด และจงให้

เหตุผลว่าค่ากลางที่ใช้ขึ้นนั้นเหมาะสมอย่างไร

ข. จงเปรียบเทียบการกระจายของข้อมูลทั้งสองชุดข้างต้น

ค. น้ำหนักของทารกแรกเกิดเบ้ซ้ายหรือเบ้ขวา มากน้อยเพียงใด

.....

.....

.....

.....

.....

11. ในการทดสอบสมมติฐานว่าระดับน้ำตาลในเลือดของผู้ป่วยกลุ่มที่ออกกำลังกายด้วยการเดินเร็วต่ำกว่าผู้ป่วยกลุ่มที่ไม่ออกกำลังกายนั้น ผู้วิจัยควรเขียนสมมติฐานทางสถิติว่าอย่างไร ถ้าต้องการทดสอบที่ระดับนัยสำคัญ .05 จากการคำนวณหาค่า Z ได้เท่ากับ -3.25 ขณะที่ค่า Z จากการเปิดค่าในตารางสำเร็จรูปที่ $\alpha .05$ เท่ากับ 1.699 และที่ .025 เท่ากับ 2.045 จะลงสรุปผลการวิจัยว่าอย่างไร

.....

.....

.....

.....