

สถิติศาสตร์ ไม่อิงพารามิเตอร์

Non-Parameter Statistics

มานะชัย รอดชื่น

คณะวิทยาศาสตร์ มหาวิทยาลัยเชียงใหม่

คำนำ

หนังสือสถิติศาสตร์ไม่อิงพารามิเตอร์เล่มนี้ได้จัดทำขึ้นสำหรับนักศึกษา นักวิจัยหรือผู้ที่สนใจสามารถใช้ประกอบการเรียนการสอนและศึกษาหาความรู้เพิ่มเติมด้วยตนเอง เพื่อให้มีความรู้และความเข้าใจในการวิเคราะห์ข้อมูลด้วยสถิติไม่อิงพารามิเตอร์ อีกทั้งมีทักษะในการเลือกใช้วิธีการทดสอบไม่อิงพารามิเตอร์ที่ถูกต้องเหมาะสมภายใต้ข้อมูลประเภทต่าง ๆ ซึ่งสามารถศึกษาประกอบกับตัวอย่างเบื้องต้นจะทำให้ผู้อ่านเข้าใจมากขึ้น ผู้อ่านหนังสือเล่มนี้ควรมีความรู้พื้นฐานทางด้านสถิติและความรู้ความเข้าใจของการใช้สถิติอิงพารามิเตอร์ ซึ่งสามารถหาอ่านได้จากตำรา และหนังสือต่าง ๆ ที่เกี่ยวข้อง

เนื้อหาในหนังสือสถิติศาสตร์ไม่อิงพารามิเตอร์เล่มนี้จะกล่าวถึงความรู้ความเข้าใจทางสถิติเบื้องต้นเพื่อนำไปสู่แนวคิดสถิติอิงพารามิเตอร์และสถิติไม่อิงพารามิเตอร์ การเลือกการทดสอบทางสถิติที่เหมาะสมในการวิจัย การทดสอบสำหรับตัวอย่างชุดเดียว การทดสอบสำหรับตัวอย่าง 2 ชุดที่สัมพันธ์กันและที่เป็นอิสระกัน การทดสอบสำหรับตัวอย่าง k ชุดที่สัมพันธ์กันและที่เป็นอิสระกัน ในการวิเคราะห์และคำนวณจะมีการนำเสนอผลลัพธ์จากโปรแกรม R ประกอบด้วยเช่นกัน

ผู้เขียนขอกราบขอบพระคุณคณาจารย์ทุกท่านที่ได้ประสิทธิ์ประสาทวิชาความรู้ทั้งทางด้านคณิตศาสตร์ และสถิติ ให้แก่ผู้เขียน ศาสตราจารย์ ดร.สำรวม จงเจริญ ที่ได้ให้แนวคิดในการเขียนหนังสือ และช่วยให้ข้อเสนอแนะเบื้องต้นในการจัดทำหนังสือ ผู้ช่วยศาสตราจารย์ นพดล เล็กสวัสดิ์ ที่กรุณาช่วยอ่าน และให้ข้อเสนอแนะในการปรับแก้ไขรายละเอียดต่าง ๆ และขอขอบคุณผู้ทรงคุณวุฒิในการประเมินหนังสือเล่มนี้ในการให้ข้อเสนอแนะอันเป็นประโยชน์ในการปรับปรุงหนังสือให้มีความสมบูรณ์มากขึ้น ผู้เขียนหวังว่าหนังสือเล่มนี้จะประโยชน์แก่นักศึกษา ครู อาจารย์ และนักวิจัย รวมทั้งผู้ที่สนใจในการวิเคราะห์ข้อมูลด้วยสถิติศาสตร์ไม่อิงพารามิเตอร์ เป็นอย่างยิ่ง หากมีข้อผิดพลาดประการใดที่เกิดขึ้นในหนังสือเล่มนี้ ผู้เขียนขออภัยมา ณ ที่นี้ และขอความกรุณาผู้อ่านแจ้งมายังผู้เขียนได้ที่ manachai.r@cmu.ac.th เพื่อจะได้นำมาปรับปรุงแก้ไขต่อไป

มานะชัย รอดชื่น

กรกฎาคม 2565

สารบัญ

คำนิยาม	iii
คำนำ	iv
สารบัญ	v
สารบัญรูป	viii
สารบัญตาราง	x

1 บทนำ

1.1 นิยามคำศัพท์	1
1.2 ประเภทของข้อมูล	4
1.3 วิธีการเก็บรวบรวมข้อมูล	5
1.4 การเลือกตัวอย่าง	5
1.5 การทดสอบสมมุติฐาน	8
1.6 การพิจารณาการปฏิเสธหรือยอมรับ H_0 จากค่า p-value และ Sig 2-tailed	11
1.7 สถิติอิงพารามิเตอร์และสถิติไม่อิงพารามิเตอร์	14
1.8 สรุปท้ายบทที่ 1	17
แบบฝึกหัดท้ายบทที่ 1	18

2 การเลือกการทดสอบทางสถิติที่เหมาะสมในการวิจัย

2.1 กำลังการทดสอบ	20
2.2 ประสิทธิภาพสัมพัทธ์	33
2.3 ประสิทธิภาพสัมพัทธ์เมื่อใกล้เคียงนัย	33
2.4 การวิเคราะห์ข้อมูลกับมาตราวัด	34
2.5 สรุปท้ายบทที่ 2	35
แบบฝึกหัดท้ายบทที่ 2	36

3

การทดสอบสำหรับตัวอย่างชุดเดียว

3.1 การทดสอบซึ่งอาศัยการแจกแจงทวินาม	37
3.2 การทดสอบควอนไทล์	48
3.3 ขีดจำกัดความคลาดเคลื่อนยินยอม	58
3.4 การทดสอบไคกำลังสอง	60
3.5 การทดสอบคอลโมโกรอฟ-สมิร์นอฟ	73
3.6 แถบความเชื่อมั่นของฟังก์ชันการแจกแจงของประชากร	83
3.7 การทดสอบของลิสส์ไฟรส์	85
3.8 การทดสอบภาวะรูปดีของคราเมอร์-ฟอนมีเชิส	87
3.9 การทดสอบซึ่งอาศัยการรันส์	90
3.10 สรุปท้ายบทที่ 3	94
แบบฝึกหัดท้ายบทที่ 3	95

4

การทดสอบสำหรับตัวอย่าง 2 ชุดที่สัมพันธ์กัน

4.1 การทดสอบของแม็กนิตาร์	100
4.2 การทดสอบด้วยเครื่องหมาย	107
4.3 การทดสอบลำดับที่โดยเครื่องหมายของวิลค็อกสัน	111
4.4 ช่วงความเชื่อมั่นของผลต่างมัธยฐาน	117
4.5 การทดสอบการสุมของฟิชเชอร์	119
4.6 สัมประสิทธิ์สหสัมพันธ์และการทดสอบนัยสำคัญ	124
4.7 สรุปท้ายบทที่ 4	135
แบบฝึกหัดท้ายบทที่ 4	136

5

การทดสอบสำหรับตัวอย่าง 2 ชุดที่เป็นอิสระกัน

5.1 การทดสอบความน่าจะเป็นแบบแม่นยำตรงของฟิชเชอร์	139
5.2 การทดสอบภาวะเอกพันธ์	144
5.3 การทดสอบมัธยฐาน	147
5.4 การทดสอบแมนน์-วิทนี	152
5.5 การทดสอบซีเกล และทูกีย์	158
5.6 การทดสอบคอลโมโกรอฟ-สมิร์นอฟ สำหรับตัวอย่าง 2 ชุด	161
5.7 การทดสอบของคราเมอร์-ฟอนมีเชิส สำหรับตัวอย่าง 2 ชุด	164

5.8 การทดสอบรันส์ของว็ลด์-วอลโฟวิทซ์ สำหรับตัวอย่าง 2 ชุด	166
5.9 การทดสอบการสุมของพีชเชอร์สำหรับตัวอย่าง 2 ชุด	169
5.10 สรุปท้ายบทที่ 5	172
แบบฝึกหัดท้ายบทที่ 5	174



การทดสอบสำหรับตัวอย่าง k ชุดที่สัมพันธ์กัน

6.1 การทดสอบค็อกแครน	180
6.2 การทดสอบฟรیدแมนต์	184
6.3 การทดสอบเดอร์บิน	186
6.4 สรุปท้ายบทที่ 6	191
แบบฝึกหัดท้ายบทที่ 6	192



การทดสอบสำหรับตัวอย่าง k ชุดที่เป็นอิสระกัน

7.1 การทดสอบไคกำลังสองสำหรับตัวอย่าง k ชุดที่เป็นอิสระกัน	197
7.2 การขยายการทดสอบมัธยฐาน	200
7.3 การทดสอบครัสคอล-วอลลิส	204
7.4 สรุปท้ายบทที่ 7	207
แบบฝึกหัดท้ายบทที่ 7	208

บรรณานุกรม	211
------------	-----

ภาคผนวก	213
---------	-----

การใช้โปรแกรม R	214
-----------------	-----

ตารางสถิติ	231
------------	-----

เฉลยแบบฝึกหัดท้ายบท	268
---------------------	-----

ดัชนี	334
-------	-----

Index	336
-------	-----

สารบัญรูป

รูปที่



หน้า

1.1	ขอบเขตของการทดสอบ $H_0 : \theta = \theta_0$ เทียบกับ $H_1 : \theta \neq \theta_0$	10
1.2	ขอบเขตของการทดสอบ $H_0 : \theta \geq \theta_0$ เทียบกับ $H_1 : \theta < \theta_0$	10
1.3	ขอบเขตของการทดสอบ $H_0 : \theta \leq \theta_0$ เทียบกับ $H_1 : \theta > \theta_0$	11
2.1	การเปรียบเทียบวิธีการใช้การทดสอบอิงพารามิเตอร์และ การทดสอบไม่อิงพารามิเตอร์	20
2.2	บริเวณวิกฤต หรือบริเวณของความผิดพลาดแบบที่ 1 กรณีทดสอบ $H_0 : \mu = \mu_0$ เทียบกับ $H_1 : \mu \neq \mu_0$	22
2.3	ความน่าจะเป็นของความผิดพลาดแบบที่ 1 และ 2 กรณีทดสอบ $H_0 : \mu = \mu_0$ เทียบกับ $H_1 : \mu \neq \mu_0$	22
2.4	ความน่าจะเป็นของความผิดพลาดแบบที่ 1 และ 2 กรณีทดสอบ $H_0 : \mu \geq \mu_0$ เทียบกับ $H_1 : \mu < \mu_0$ เมื่อ $\bar{x}_a = \mu_0 - z_{1-\alpha} \frac{\sigma}{\sqrt{n}}$	23
2.5	ความน่าจะเป็นของความผิดพลาดแบบที่ 1 และ 2 กรณีทดสอบ $H_0 : \mu \leq \mu_0$ เทียบกับ $H_1 : \mu > \mu_0$ เมื่อ $\bar{x}_b = \mu_0 + z_{1-\alpha} \frac{\sigma}{\sqrt{n}}$	23
2.6	ความสัมพันธ์ระหว่าง α และ β กรณีทดสอบ $H_0 : \mu \leq \mu_0$ เทียบกับ $H_1 : \mu > \mu_0$ เมื่อค่า α เพิ่มมากขึ้น	23
2.7	ความสัมพันธ์ระหว่าง α และ β กรณีทดสอบ $H_0 : \mu \leq \mu_0$ เทียบกับ $H_1 : \mu > \mu_0$ เมื่อค่า α ลดลง	23
2.8	ความสัมพันธ์ระหว่าง α และ β กรณีทดสอบ $H_0 : \mu = \mu_0$ เทียบกับ $H_1 : \mu \neq \mu_0$ เมื่อค่า μ_0 และ μ_1 มีค่าต่างกันมาก ๆ	24
2.9	เส้นโค้งลักษณะเฉพาะดำเนินการ และกำลังการทดสอบ	31
2.10	เส้นโค้งของความน่าจะเป็นของความผิดพลาดแบบที่ 2 เมื่อ $n = 16, 25, 100$	32
2.11	เส้นโค้งของกำลังการทดสอบ เมื่อ $n = 16, 25, 100$	32
2.12	การเปรียบเทียบค่า β กรณีการทดสอบ $H_0 : \mu = \mu_0$ เทียบกับ $H_1 : \mu \neq \mu_0$ และการทดสอบ $H_0 : \mu \leq \mu_0$ เทียบกับ $H_1 : \mu > \mu_0$	33



3.1	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P = P_0$ เทียบกับ $H_1 : P \neq P_0$	39
3.2	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P \geq P_0$ เทียบกับ $H_1 : P < P_0$	39
3.3	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P \leq P_0$ เทียบกับ $H_1 : P > P_0$	40
3.4	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P(X \leq x_0) \geq p$ เทียบกับ $H_1 : P(X \leq x_0) < p$	49
3.5	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P(X < x_0) \leq p$ เทียบกับ $H_1 : P(X < x_0) > p$	50
3.6	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P(X \leq x_0) \geq p$ และ $P(X < x_0) \leq p$ เทียบกับ $H_1 : P(X \leq x_0) < p$ หรือ $P(X < x_0) > p$	51
3.7	บริเวณวิกฤตการณ์ทดสอบ $H_0 : P_{ij} = P_i \cdot P_j$ เทียบกับ $H_1 : P_{ij} = P_i \cdot P_j \exists i, j$	71
3.8	ค่าของ $F_0(x)$ และ $F_n(x)$ ในกรณีที่ $F_0(x)$ เป็นฟังก์ชันการแจกแจง ของตัวแปรสุ่มที่ต่อเนื่อง	75
3.9	ค่าของ $F_0(x)$ และ $F_n(x)$ ในกรณีที่ $F_0(x)$ เป็นฟังก์ชันการแจกแจง ของตัวแปรสุ่มที่ไม่ต่อเนื่อง	75
3.10	แถบความเชื่อมั่นของ $F(x)$ จากตัวอย่างที่ 3.24 ที่ช่วงความเชื่อมั่น 90%	84

สารบัญตาราง

ตารางที่



หน้า

1.1	การตัดสินใจในการทดสอบสมมุติฐาน	9
1.2	การทดสอบอิงพารามิเตอร์และไม่อิงพารามิเตอร์จำแนกตามลักษณะ มาตราวัดของข้อมูล	16
2.1	ค่าความน่าจะเป็นของความผิดพลาดแบบที่ 2 (β) และ กำลังการทดสอบ ($1-\beta$) ของตัวอย่างที่ 2.6	30
2.2	ค่าความน่าจะเป็นของความผิดพลาดแบบที่ 2 (β) และ กำลังการทดสอบ ($1-\beta$) ของตัวอย่างที่ 2.6 กรณีที่ $n = 25,100$	31
2.3	ค่า ARE ของการเปรียบเทียบการทดสอบอิงพารามิเตอร์กับ การทดสอบไม่อิงพารามิเตอร์	34
2.4	การวิเคราะห์ข้อมูลเบื้องต้น และการทดสอบสมมุติฐาน จำแนกตามมาตราวัด	35
3.1	ตารางการจรของตัวแปร 2 ตัวแปร	70
3.2	ค่าวิกฤตของการทดสอบของคราเมอร์ฟอนมีเชิส	88
4.1	ตารางการจรของตัวแปร 2 ตัวแปร ขนาด 2×2	100
4.2	บริเวณวิกฤตการทดสอบด้วยเครื่องหมาย	109
4.3	บริเวณวิกฤตการทดสอบลำดับที่โดยเครื่องหมายของวิลค็อกซัน	113
4.4	แสดงค่าเฉลี่ยของข้อมูลทุกคู่ที่เป็นไปได้ จากตัวอย่างที่ 4.8	118
4.5	บริเวณวิกฤตการทดสอบของสเปียร์แมน	126
5.1	ตารางการจรของข้อมูลการทดสอบของฟิชเชอร์	140
5.2	ลักษณะของข้อมูลสำหรับการทดสอบภาวะเอกพันธ์ ของสัดส่วนของประชากร	145
5.3	บริเวณวิกฤตการทดสอบภาวะเอกพันธ์ของสัดส่วนของ 2 ประชากร	145
5.4	ตารางการจรขนาด 2×2 สำหรับการทดสอบภาวะเอกพันธ์ ของสัดส่วนของ 2 ประชากร	146
5.5	ตารางการจรขนาด 2×2 สำหรับการทดสอบมัธยฐาน	148



5.6	บริเวณวิกฤตการทดสอบแมนน์-วิทนี	154
5.7	แสดงผลต่างของข้อมูลทุกคู่ที่เป็นไปได้ จากตัวอย่างที่ 5.5	157
5.8	บริเวณวิกฤตการทดสอบแปรปรวน 2 ประชากร โดยใช้การทดสอบแมนน์-วิทนี	158
6.1	ข้อมูลการทดสอบของค็อกแครน	180
6.2	ข้อมูลการทดสอบฟรีดแมนต์	184
7.1	ข้อมูลการทดสอบความเหมือนกันของสัดส่วนของ k ประชากร	198
7.2	ข้อมูลการทดสอบความเหมือนกันของมัธยฐานของ k ประชากร	201
7.3	ข้อมูลของการทดสอบของคริสต์คอล-วอลลิส	204
A1	ความน่าจะเป็นสะสมของการแจกแจงทวินาม $P(X \leq c) = \sum_{x=0}^c \binom{n}{x} p^x (1-p)^{n-x}$	231
A2	ค่าของขีดจำกัดล่าง $\hat{P}_L$ และขีดจำกัดบน $\hat{P}_U$ เมื่อทราบค่าของ n , $(1-\alpha)$ (จำนวนของตัวอย่างสุ่มที่เกิดเหตุการณ์ที่สนใจ) เมื่อ $n = 1, 2, 10$	241
A3	ความน่าจะเป็นสะสมของการแจกแจงปัวซอง $P(X \leq c) = \sum_{x=0}^c \frac{e^{-\lambda} \lambda^x}{x!}$	245
A4	ความน่าจะเป็นสะสมของการแจกแจงปกติมาตรฐาน $F(x) = P(Z < x) = \int_{-\infty}^x \frac{e^{-t^2/2}}{\sqrt{2\pi}} dt$	250
A5	ความน่าจะเป็นสะสมของการแจกแจงที่ $H(T < t_{p,v}) = p$	251
A6	ความน่าจะเป็นสะสมของการแจกแจงไคกำลังสอง $P(\chi^2 \leq \chi_{p,v}^2) = p$	252
A7	ค่าควอนไทล์การทดสอบซึ่งอาศัยรันส์ของวัลด์-วอลโฟวิทซ์	254
A8	ขนาดตัวอย่างสำหรับขีดจำกัดความคลาดเคลื่อนยินยอมแบบข้างเดียว	255



A9	ขนาดตัวอย่างสำหรับขีดจำกัดความคลาดเคลื่อนยินยอมแบบสองข้าง เมื่อ $r + m = 2$	255
A10	ค่าควอนไทล์การทดสอบลำดับที่โดยเครื่องหมายของวิลค็อกซัน	256
A11	ค่าควอนไทล์การทดสอบแมนน์-วิทนีย์	257
A12	ค่าควอนไทล์การทดสอบของสเปียร์แมน	260
A13	ค่าควอนไทล์การทดสอบของเคนดัล	261
A14	ค่าควอนไทล์การทดสอบของคอลโมโกรอฟ	262
A15	ค่าควอนไทล์การทดสอบของลิลลีโฟร์ส	263
A16	ค่าควอนไทล์การทดสอบของสมิร์นอฟ สำหรับตัวอย่างสองชุด ที่มีขนาดเท่ากัน	264
A17	ค่าควอนไทล์การทดสอบของสมิร์นอฟ สำหรับตัวอย่างสองชุด ที่มีขนาดไม่เท่ากัน	265
A18	ค่าควอนไทล์การทดสอบด้วยเครื่องหมายแบบสองด้าน (จากการแจกแจงทวินาม เมื่อ $p = 0.5$)	267

หนังสือสถิติศาสตร์ไม่อิงพารามิเตอร์ (Non-Parameter Statistics)

หนังสือสถิติศาสตร์ไม่อิงพารามิเตอร์ เล่มนี้ได้จัดทำขึ้นจากประสบการณ์การสอน และวิจัยเกี่ยวกับวิชาสถิติศาสตร์ไม่อิงพารามิเตอร์ ในระดับปริญญาตรี และปริญญาโท ผู้อ่านหนังสือเล่มนี้ควรมีความรู้พื้นฐานทางด้านสถิติ และความรู้ความเข้าใจของ ข้อตกลงเบื้องต้นของการใช้สถิติอิงพารามิเตอร์ จะทำให้ผู้อ่านเข้าใจได้ง่ายขึ้น โดย เนื้อหาของหนังสือเล่มนี้จะกล่าวถึงหลักการและตัวอย่างการใช้สถิติไม่อิงพารามิเตอร์ พร้อมทั้งแสดงตัวอย่างการคำนวณ และผลลัพธ์จากโปรแกรม R เพื่อให้ผู้อ่านมีความ เข้าใจและสามารถนำไปใช้ในการวิเคราะห์ข้อมูลต่าง ๆ ได้ตามความเหมาะสมกับตัวแปร ที่ศึกษานั้น ๆ อีกทั้งยังมีแบบฝึกหัด พร้อมเฉลยเพื่อให้ผู้อ่านมีความเข้าใจมากขึ้น



CHIANG MAI
UNIVERSITY PRESS

ISBN 978-616-398-748-8



9 786163 987488



บทนำ

วิชาสถิติ ได้มีการแบ่งเนื้อหาออกเป็นสองส่วน คือ สถิติเชิงพรรณนา (Descriptive Statistics) และสถิติเชิงอนุมาน (Inferential Statistics) สถิติเชิงพรรณนาคือกระบวนการทางสถิติที่ใช้อธิบายหรือสรุปรายละเอียดของข้อมูลในรูปแบบของตัวเลข หรือแผนภาพต่าง ๆ โดยไม่ใช้วิธีการเชิงความน่าจะเป็นเข้ามาเกี่ยวข้อง ส่วนสถิติเชิงอนุมานเป็นกระบวนการที่ใช้ข้อมูลจากตัวอย่างสุ่ม (Random Sample) มาหาข้อสรุปเกี่ยวกับคุณลักษณะของประชากรที่เรียกว่าพารามิเตอร์ (Parameter) สถิติเชิงอนุมานประกอบด้วย การประมาณค่า (Estimation) และการทดสอบสมมติฐาน (Testing of Hypothesis) ซึ่งเป็นกระบวนการที่สามารถประเมินความถูกต้องภายใต้ความน่าจะเป็นที่เกี่ยวข้อง

ในบทนี้ จะขอทบทวนคำศัพท์ต่าง ๆ ทางสถิติที่เกี่ยวข้อง ประเภทของข้อมูล วิธีการเก็บรวบรวมข้อมูล การเลือกตัวอย่างแบบใช้และไม่ใช้ความน่าจะเป็น การทดสอบสมมติฐาน การพิจารณาความน่าจะเป็นของสถิติทดสอบในการปฏิเสธสมมติฐาน สถิติอิงพารามิเตอร์และสถิติไม่อิงพารามิเตอร์ เพื่อให้ผู้อ่านได้ทบทวนเนื้อหาพื้นฐานทางสถิติ

1.1 นิยามคำศัพท์

สถิติศาสตร์ (Statistics) หมายถึง ศาสตร์ว่าด้วยกระบวนการที่ใช้ในการวางแผนเก็บรวบรวมข้อมูล การจัดการข้อมูล การวิเคราะห์ การสรุปผล และแปลผลจากข้อมูล เพื่อหาข้อสรุปในการตัดสินใจแก้ปัญหา รวมทั้งการนำเสนอข้อมูลด้วยวิธีเชิงสถิติ

สถิติเชิงพรรณนา (Descriptive Statistics) คือ กระบวนการต่าง ๆ เชิงสถิติในการสรุปลักษณะสำคัญของข้อมูลที่ได้เก็บรวบรวมมา เพื่ออธิบายลักษณะหรือสภาพของข้อมูลซึ่งไม่ใช้วิธีการเชิงความน่าจะเป็น นั่นคือ เป็นการนำเสนอข้อมูล การวิเคราะห์ข้อมูลเบื้องต้น เช่น การนำเสนอข้อมูลสรุปในเชิงตัวเลข (เช่น ร้อยละ ค่าต่ำสุด ค่าสูงสุด) หรือการสรุปในรูปของแผนภาพต่าง ๆ (เช่น ฮิสโทแกรม แผนภูมิรูปวงกลม) การวัดแนวโน้มเข้าสู่ส่วนกลาง (เช่น ค่าเฉลี่ย มัธยฐาน ฐานนิยม) การวัดการกระจายของข้อมูล (เช่น ส่วนเบี่ยงเบนมาตรฐาน) เป็นต้น

สถิติเชิงอนุมาน (Inferential Statistics) เป็นกระบวนการที่วาดด้วยการหาข้อสรุปเกี่ยวกับคุณลักษณะของประชากร โดยอาศัยข้อมูลจากตัวอย่างสุ่ม ประกอบไปด้วยการประมาณค่า และการทดสอบสมมติฐาน

ประชากร (Population) คือ เซตหรือกลุ่มของหน่วยต่าง ๆ ทั้งหมด ในขอบข่ายที่สนใจ ทำการศึกษา ซึ่งแบ่งได้เป็น

- ประชากรอันตะ หรือประชากรจำกัด (Finite Population) คือ เซตของหน่วยต่าง ๆ ทั้งหมด ในขอบข่ายที่สนใจศึกษา ซึ่งมีจำนวนสมาชิกที่แน่นอน หรือมีสมาชิกเป็นจำนวนจำกัด โดยทั่วไปจะใช้ N แทนขนาดของประชากร

- ประชากรไม่จำกัด หรือประชากรอนันต์ (Infinite Population) คือ เซตของหน่วยต่าง ๆ ในขอบข่ายที่สนใจศึกษา ที่ไม่ทราบจำนวนที่แน่นอน

ตัวอย่าง (Sample) คือ เซตหรือกลุ่มของหน่วยต่าง ๆ บางส่วนของประชากรที่ถูกเลือกมาเป็นตัวแทนของประชากร

กรอบตัวอย่าง (Sampling Frame) คือ สิ่งที่ใช้แสดงถึงหน่วยทุกหน่วยในประชากรที่ต้องการเลือกตัวอย่าง ซึ่งสามารถระบุหน่วยทุกหน่วยของประชากรที่ศึกษาได้ครบถ้วนและสมบูรณ์ เช่น บัญชีรายชื่อได้แก่ รายชื่อหน่วย พร้อมตำบลที่อยู่ หรือแผนที่ที่คลุมบริเวณที่ศึกษา พร้อมการแบ่งเป็นหน่วยพื้นที่ กรอบตัวอย่างที่ดีจะต้องแสดงตำแหน่งที่อยู่หน่วยต่าง ๆ และรายชื่อจะต้องไม่ซ้ำซ้อน มีความถูกต้องและทันสมัย อีกทั้งไม่มีบัญชีรายชื่อของประชากรอื่น ๆ ปะปนอยู่

ตัวอย่างสุ่ม (Random Sample) คือ เซตของตัวแปรสุ่มที่สนใจศึกษา ที่ถูกสุ่มมาจากประชากร ขนาดตัวอย่างสุ่ม เขียนแทนด้วย n โดยทั่วไปเมื่อกล่าวถึงตัวอย่างสุ่มจะหมายถึง ตัวแปรสุ่มที่แต่ละตัวมีการแจกแจงเหมือนกัน และเป็นอิสระซึ่งกันและกัน

หน่วยตัวอย่าง (Sampling Unit) คือ หน่วยที่เป็นสมาชิกของประชากรซึ่งอาจถูกเลือกเข้ามาเป็นหน่วยตัวอย่าง หน่วยตัวอย่างอาจเป็นหน่วยหรือกลุ่มของหน่วยที่ให้ข้อมูล

แผนแบบการทดลอง (Experimental Design) เป็นกระบวนการเกี่ยวกับการวางแผนการทดลองและเก็บรวบรวมข้อมูลเพื่อหาข้อสรุป ข้อเท็จจริงใหม่ ๆ โดยใช้การวิเคราะห์เชิงสถิติ ผู้ทดลองสามารถควบคุมกระบวนการต่าง ๆ ในการทดลองให้มีความคลาดเคลื่อนน้อยที่สุด ซึ่งผลจากการทดลองนั้นอาจจะได้ข้อสรุปใหม่ หรือข้อสรุปที่เป็นการยืนยัน หรือขัดแย้งกับการทดลองที่ผ่านมา

ทรีตเมนต์ (Treatment) เป็นสิ่งเร้าหรือตัวกระตุ้นที่กำหนดให้กับหน่วยทดลองเพื่อสังเกตผล หรือใช้เปรียบเทียบผลระหว่างกัน หรือเป็นระดับของปัจจัยที่ศึกษา

หน่วยทดลอง (Experimental Unit) คือหน่วยต่าง ๆ ที่ใช้สำหรับทำการทดลองกับทรีตเมนต์ 1 ทรีตเมนต์ หน่วยทดลองอาจจะเป็น คน สัตว์ สิ่งของ ต้นไม้ แปลงทดลอง (ทางด้านเกษตรกรรม)

การทดลองสุ่ม (Random Experiment) เป็นการทดลองที่ไม่ทราบผลลัพธ์ที่จะเกิดขึ้นแน่นอน โดยจะมีผลลัพธ์ที่เป็นไปได้มากกว่า 1 ผลลัพธ์

ปัจจัย (Factor) สิ่งที่น่าสนใจศึกษาและคาดว่าจะมีผลต่อหน่วยทดลอง นั่นคือ ตัวแปรอิสระที่สามารถควบคุมได้และเป็นสาเหตุหนึ่งของความแปรผันที่เกิดขึ้น ในแต่ละปัจจัยอาจจะแบ่งได้หลายระดับ แต่ละระดับเรียกว่า ทรีตเมนต์

การทำซ้ำ (Replication) การที่แต่ละทรีตเมนต์ถูกกำหนดให้กับหน่วยทดลองอย่างสุ่มมากกว่าหนึ่งครั้งขึ้นไป หรือปรากฏในหน่วยทดลองมากกว่า 1 หน่วยทดลอง ประโยชน์ของการทำซ้ำนั้น เพื่อช่วยในการประมาณค่าความคลาดเคลื่อนจากการทดลอง ทำให้การทดลองมีความแม่นยำมากขึ้น

แผนแบบสุ่มสมบูรณ์ (Completely Randomized Design: CRD) เป็นแผนการทดลองที่กำหนดทรีตเมนต์ให้กับหน่วยทดลองทั้งหมดโดยการสุ่ม โดยที่หน่วยทดลองแต่ละหน่วยมีโอกาสที่จะได้รับทรีตเมนต์ใดทรีตเมนต์หนึ่งเท่า ๆ กัน

แผนแบบบล็อกสมบูรณ์เชิงสุ่ม (Randomized Complete Block Design: RCBD) ในกรณีที่หน่วยทดลองแต่ละหน่วยมีลักษณะที่แตกต่างกัน (Heterogeneous) ดังนั้นจะจัดกลุ่มให้กับหน่วยทดลอง โดยที่กำหนดให้หน่วยทดลองที่อยู่ภายในกลุ่มเดียวกันมีลักษณะที่คล้ายคลึงกัน ระหว่างกลุ่มมีลักษณะที่แตกต่างกัน กลุ่มของหน่วยทดลอง จะเรียกว่า บล็อก (Block)

แผนแบบบล็อกสมบูรณ์เชิงสุ่มเป็นแผนแบบการทดลองที่กำหนดทรีตเมนต์ให้หน่วยทดลองในบล็อกแต่ละบล็อกอย่างสุ่ม โดยทรีตเมนต์ต้องปรากฏครบทุกทรีตเมนต์เป็นจำนวนเท่ากันในแต่ละบล็อก โดยทั่วไปทรีตเมนต์จะปรากฏในแต่ละบล็อกเพียง 1 ครั้งเท่านั้น

ปริภูมิตัวอย่าง (Sample Space) คือ เซตของผลลัพธ์ที่เป็นไปได้ทั้งหมดของการทดลองสุ่ม แทนด้วย S

ตัวแปรสุ่ม (Random Variable) คือ ฟังก์ชันที่มีค่าเป็นจำนวนจริง ซึ่งกำหนดให้กับแต่ละสมาชิกในปริภูมิตัวอย่าง

ชนิดของตัวแปรสุ่ม

- 1) ตัวแปรสุ่มไม่ต่อเนื่อง (Discrete Random Variable) คือ ตัวแปรสุ่มซึ่งมีค่าเป็นจำนวนที่สามารถนับได้ เช่น จำนวนเหรียญที่ออกหัวจากการทดลอง จำนวนนักศึกษาที่สอบผ่าน
- 2) ตัวแปรสุ่มต่อเนื่อง (Continuous Random Variable) คือ ตัวแปรสุ่มที่มีค่าอยู่บนช่วงของจำนวนจริง เช่น น้ำหนัก ส่วนสูง อุณหภูมิ

พารามิเตอร์ (Parameter) คือ ค่าคงตัวที่แสดงคุณลักษณะบางประการของประชากร อาจจะเป็นค่าคงตัวส่วนหนึ่งที่ใช้ในการอธิบายฟังก์ชันมวลความน่าจะเป็น หรือฟังก์ชันความหนาแน่นความน่าจะเป็นของวงศ์การแจกแจง เช่น การแจกแจงปกติมีพารามิเตอร์คือ ค่าเฉลี่ย μ , ค่าความแปรปรวน σ^2

สถิติ (Statistic) มี 2 ชนิด ได้แก่

ตัวสถิติ คือ ฟังก์ชันของตัวแปรสุ่มในตัวอย่างสุ่มที่ใช้ในการอนุมานค่าพารามิเตอร์

ค่าสถิติ คือ ค่าของตัวสถิติที่คำนวณได้จากข้อมูลตัวอย่าง

1.2 ประเภทของข้อมูล

การแบ่งประเภทของข้อมูลมีหลายวิธี ดังนี้

1. แบ่งตามลักษณะของข้อมูลเป็น 2 ประเภท คือ

1.1 ข้อมูลเชิงคุณภาพ (Qualitative Data) เป็นข้อมูลที่ไม่สามารถวัดค่าเป็นตัวเลขได้ เป็นข้อมูลที่แสดงลักษณะ ประเภท หรือสมบัติในเชิงคุณภาพต่าง ๆ ไม่สามารถบอกขนาดความแตกต่างได้ว่ามากหรือน้อยกว่ากันแค่ไหน เช่น เพศ ระดับการศึกษา อาชีพ

1.2 ข้อมูลเชิงปริมาณ (Quantitative Data) เป็นข้อมูลที่สามารถวัดหรือการนับค่าแสดงเป็นตัวเลขได้ ใช้แทนขนาดหรือปริมาณ บอกความแตกต่าง หรือระบุค่ามากหรือน้อยได้ เช่น รายได้ ส่วนสูง อายุ น้ำหนัก จำนวนผู้มาใช้บริการตู้เอทีเอ็ม แบ่งได้ 2 ชนิด คือ

1.2.1 ข้อมูลไม่ต่อเนื่อง (Discrete Data) หมายถึง ข้อมูลที่มีค่าเป็นจำนวนที่สามารถนับได้ เช่น จำนวนสิ่งของ จำนวนคน

1.2.2 ข้อมูลต่อเนื่อง (Continuous Data) หมายถึง ข้อมูลที่มีค่าเป็นได้ทุกค่าของจำนวนจริง เช่น น้ำหนัก รายได้ ส่วนสูง

2. แบ่งตามแหล่งที่มาของข้อมูล มี 2 ประเภท คือ

2.1 ข้อมูลปฐมภูมิ (Primary Data) คือ ข้อมูลที่ผู้ใช้ข้อมูลเป็นผู้จัดเก็บรวบรวมจากเจ้าของหรือต้นกำเนิดของข้อมูลโดยตรง อาจเก็บโดยการสัมภาษณ์ ทดลอง หรือสังเกตการณ์ ซึ่งอาจจะมีรายละเอียดครบตามกำหนดที่ผู้ใช้ต้องการ แต่จะเสียเวลาและค่าใช้จ่ายมาก

2.2 ข้อมูลทุติยภูมิ (Secondary Data) คือ ข้อมูลที่ผู้ใช้ไม่ได้เก็บรวบรวมจากหน่วยที่ให้ข้อมูลโดยตรง มีผู้เก็บรวบรวมจากหน่วยที่ให้ข้อมูลไว้แล้ว การใช้ข้อมูลทุติยภูมิทำให้ประหยัดเวลาและค่าใช้จ่าย แต่อาจจะไม่ตรงกับวัตถุประสงค์ หรือมีรายละเอียดไม่เพียงพอกับผู้ใช้ข้อมูล และผู้ใช้ข้อมูลไม่ทราบถึงข้อผิดพลาดของข้อมูล อาจทำให้ข้อสรุปผิดพลาด

3. แบ่งตามมาตราวัดของข้อมูล (Measurement Scale) แบ่งได้ 4 มาตรา ดังนี้

3.1 มาตรานามบัญญัติ (Nominal Scale) เป็นมาตราวัดที่ใช้กับข้อมูลเชิงคุณภาพ โดยแบ่งข้อมูลเป็นกลุ่ม เช่น เพศ อาชีพ สถานภาพสมรส กลุ่มเลือด

3.2 มาตราอันดับ (Ordinal Scale) เป็นมาตราวัดของข้อมูลที่สามารถบอกลักษณะเชิงเปรียบเทียบตามอันดับความสำคัญ แสดงความแตกต่างโดยอาศัยเกณฑ์ใดเกณฑ์หนึ่ง แต่ไม่สามารถบอกปริมาณความแตกต่างที่แท้จริงได้ เช่น ระดับการศึกษา เกรดของวิชาสถิติ

3.3 มาตราช่วง (Interval Scale) เป็นมาตราวัดข้อมูลเชิงปริมาณที่มีมาตรฐานการวัดที่แน่นอน สามารถหาความแตกต่างที่แท้จริงในแต่ละช่วงได้ แต่ไม่สามารถเปรียบเทียบอัตราความแตกต่างเป็นจำนวนเท่าได้ นั่นคือ ค่าเริ่มต้นที่ 0 มีความหมาย เช่น คะแนนจากการทดสอบวัดความรู้ การวัดอุณหภูมิ เป็นองศาเซลเซียส หรือองศาฟาเรนไฮต์ สามารถสรุปได้ว่า ค่า 0°C ไม่ได้หมายถึงไม่มีอุณหภูมิ หรือผลต่างระหว่าง 50°F กับ 0°F เป็น 2 เท่าของผลต่างระหว่าง 25°F กับ 0°F แต่ไม่สามารถบอกได้ว่า ความร้อน 50°F ร้อนเป็น 2 เท่าของ 25°F เพราะเมื่อเปลี่ยนเป็น

องศาเซลเซียสแล้วจะได้ว่า $50^{\circ}\text{F} = 10^{\circ}\text{C}$ และ $25^{\circ}\text{F} = -3.9^{\circ}\text{C}$ ซึ่งไม่เป็นอัตราส่วน 2:1 ($0^{\circ}\text{C} = 32^{\circ}\text{F}$, $100^{\circ}\text{C} = 212^{\circ}\text{F}$)

3.4 มาตราอัตราส่วน (Ratio Scale) เป็นมาตราวัดที่สมบูรณ์มากที่สุด จุดเริ่มต้น คือ จุดศูนย์ที่แท้จริง สามารถบอกขนาดความแตกต่าง และสามารถเปรียบเทียบความแตกต่างได้ เช่น ระยะทาง ความสูง รายได้ น้ำหนัก เป็นต้น

1.3 วิธีการเก็บรวบรวมข้อมูล

การเก็บรวบรวมข้อมูลอาจใช้วิธีการดังต่อไปนี้

1. **สำมะโน (Census)** คือ การเก็บรวบรวมข้อมูลที่สนใจศึกษาจากทุกหน่วยในประชากร เช่น การสำมะโนประชากรของสำนักงานสถิติแห่งชาติ

2. **การสำรวจตัวอย่าง (Sample Survey)** คือ การเก็บรวบรวมข้อมูลจากกลุ่มย่อยหรือบางหน่วยของประชากร เช่น สนใจศึกษาค่าใช้จ่ายต่อวันของนักศึกษาในมหาวิทยาลัยแห่งหนึ่ง จึงได้เก็บข้อมูลจากนักศึกษาของมหาวิทยาลัยแห่งนี้มา 200 คน (นักศึกษา 200 คน คือตัวอย่าง)

3. **การทดลอง (Experiment)** คือ กระบวนการเก็บข้อมูลที่ได้จากการวางแผนการทดลอง

1.4 การเลือกตัวอย่าง

การเลือกตัวอย่างแบ่งเป็น 2 ประเภท คือ

1. **การเลือกตัวอย่างแบบใช้ความน่าจะเป็น (Probability Sampling)** เป็นวิธีการเลือกตัวอย่างที่ใช้ความน่าจะเป็นในการเลือกหน่วยตัวอย่างในขั้นตอนต่าง ๆ ซึ่งทำให้สามารถประมาณคุณภาพของการอนุมานคุณลักษณะของประชากรที่สนใจศึกษาได้ เช่น

1.1 การเลือกตัวอย่างสุ่มแบบง่าย (Simple Random Sampling)

เป็นการเลือกหน่วยตัวอย่างขนาด n จากประชากรขนาด N โดยที่หน่วยตัวอย่างแต่ละหน่วยในประชากรมีโอกาสถูกเลือกเข้ามาในตัวอย่างเท่า ๆ กัน วิธีการเลือกหน่วยตัวอย่าง โดยส่วนมากใช้การจับสลาก ตารางเลขสุ่ม (Random Number) หรือใช้เลขสุ่มจากโปรแกรมสำเร็จรูป เป็นต้น

การเลือกตัวอย่างแบบคืนที่ (Sampling with Replacement) เป็นการเลือกตัวอย่างขนาด n ที่หน่วยตัวอย่างแต่ละหน่วยมีโอกาสถูกเลือกเท่า ๆ กัน ในการเลือกแต่ละครั้ง และทำการคืนหน่วยตัวอย่างที่เลือกแล้วกลับเข้าไปในประชากรก่อนเลือกหน่วยตัวอย่างถัดไป ทำให้มีโอกาสได้หน่วยตัวอย่างเดียวกันมากกว่า 1 ครั้ง และความน่าจะเป็นของการเลือกได้ตัวอย่างมีค่าเท่ากับ $\frac{1}{N^n}$ ส่วนการเลือกตัวอย่างแบบไม่คืนที่ (Sampling without Replacement) เป็นการเลือกตัวอย่างโดยไม่คืนหน่วยตัวอย่างที่เลือกกลับไปในประชากรก่อนเลือกหน่วยตัวอย่างถัดไป ทำให้ไม่มีโอกาสได้หน่วยตัวอย่างเดียวกันมากกว่า 1 ครั้ง ความน่าจะเป็นของการเลือกได้ตัวอย่างมีค่าเท่ากับ $\frac{1}{N \cdot C_n}$

1.2 การเลือกตัวอย่างสุ่มแบบมีระบบ (Systematic Random Sampling)

เป็นวิธีการเลือกหน่วยตัวอย่าง โดยทุก ๆ k หน่วยตัวอย่างของประชากรจะเลือกหน่วยตัวอย่างได้ 1 หน่วย เป็นวิธีการสุ่มที่ทำได้ง่าย และสะดวก โดยเฉพาะกรณีที่ไม่มีการรอบตัวอย่าง เช่น ต้องการประมาณค่าใช้จ่ายของนักท่องเที่ยวในการซื้อสินค้าของร้านค้าแห่งหนึ่ง หากไม่ทราบจำนวนนักท่องเที่ยวที่มาซื้อสินค้าทั้งหมด จึงกำหนดวิธีการสุ่มโดยกำหนดว่าจะสำรวจค่าใช้จ่ายในการซื้อสินค้าของนักท่องเที่ยวทุก ๆ 10 คน จนกระทั่งได้นักท่องเที่ยวครบ 50 คน

ถ้าต้องการเลือกหน่วยตัวอย่างขนาด n จากประชากร N จะทำโดยกำหนดหมายเลขแก่ประชากร จากนั้นทำการหาช่วง (Interval : I) โดย $I = \frac{N}{n} \approx k$ แล้วทำการเลือกเลขสุ่มหมายเลข 1 ถึง k มาหนึ่งหมายเลขอย่างง่ายด้วยความน่าจะเป็นที่เท่า ๆ กัน มีรูปแบบการสุ่ม 2 แบบ คือ การเลือกตัวอย่างสุ่มแบบมีระบบเส้นตรง (Linear Systematic Random Sampling) ในกรณีที่ I เป็นจำนวนเต็ม และการเลือกตัวอย่างสุ่มแบบมีระบบวงกลม (Circular Systematic Random Sampling) ในกรณีที่ I ไม่เป็นจำนวนเต็ม

1.3 การเลือกตัวอย่างสุ่มแบบแบ่งชั้นภูมิ (Stratified Random Sampling)

กรณีที่หน่วยต่าง ๆ ของประชากรมีลักษณะแตกต่างกันแต่สามารถจัดเป็นหมวดหมู่ได้ กรณีนี้ทำการแบ่งหน่วยต่าง ๆ ของประชากรออกเป็นส่วยย่อย เรียกว่า ชั้นภูมิ (Stratum) โดยที่หน่วยตัวอย่างในชั้นภูมิเดียวกันมีลักษณะคล้ายคลึงกันมากที่สุด (Homogeneous) และระหว่างชั้นภูมิมีลักษณะแตกต่างกันมากที่สุด (Heterogeneous) ภายใต้คุณลักษณะที่ศึกษา จะเลือกตัวอย่างจากแต่ละชั้นภูมิ โดยเลือกตัวอย่างสุ่มแบบง่าย หรือเลือกตัวอย่างแบบมีระบบ หรือการเลือกตัวอย่างโดยใช้ความน่าจะเป็นในการเลือกตัวอย่างเชิงสัดส่วนกับขนาดตัวอย่าง (Sampling with Probability Proportion to Size)

1.4 การเลือกตัวอย่างสุ่มแบบกลุ่ม (Cluster Random Sampling)

หากคุณลักษณะที่สนใจของประชากรมีหลายแบบ และประชากรมีขนาดใหญ่ การกำหนดกรอบตัวอย่างอาจทำได้ไม่ครบถ้วน ดังนั้น กรณีเช่นนี้จะใช้วิธีแบ่งประชากรออกเป็นกลุ่ม ๆ เรียกว่า คลัสเตอร์ (Cluster) โดยให้ประชากรที่อยู่ในกลุ่มเดียวกันมีลักษณะที่แตกต่างกัน (Heterogeneous) และประชากรที่อยู่ระหว่างกลุ่มมีลักษณะที่คล้ายคลึงกัน (Homogeneous) แล้วทำการสุ่มคลัสเตอร์จำนวนหนึ่งมาเป็นตัวแทนของประชากร โดยใช้วิธีการเลือกตัวอย่างสุ่มแบบง่าย หรือการเลือกตัวอย่างแบบมีระบบก็ได้ แล้วเก็บรวบรวมข้อมูลจากคลัสเตอร์ที่ตกเป็นหน่วยตัวอย่าง วิธีการนี้เป็นการเลือกตัวอย่างสุ่มแบบกลุ่มชั้นเดียว

วิธีการเลือกตัวอย่างสุ่มแบบกลุ่มสองชั้น คือ หลังจากแบ่งประชากรออกเป็นกลุ่ม ๆ (คลัสเตอร์) แล้ว ขั้นที่ 1 ทำการสุ่มคลัสเตอร์มา ขั้นที่ 2 สุ่มตัวอย่างจากคลัสเตอร์ที่สุ่มมาจากขั้นตอนที่ 1 ซึ่งตัวอย่างจากขั้นที่ 2 นี้ เป็นตัวอย่างขั้นสุดท้าย

2. การเลือกตัวอย่างแบบไม่ใช้ความน่าจะเป็น (Non-Probability Sampling) เป็นการเลือกตัวอย่างโดยไม่ใช้ความน่าจะเป็นในการเลือกหน่วยตัวอย่างในขั้นตอนต่าง ๆ (เป็นการเลือก

หน่วยตัวอย่างโดยไม่เป็นเชิงสุ่ม) การเลือกโดยวิธีนี้จะทำให้ได้ตัวอย่างเอนเอียง (Biased Sample) เกิดความคลาดเคลื่อนเชิงระบบ (Systematic Error) ซึ่งจะแตกต่างจากความคลาดเคลื่อนเชิงสุ่ม (Random Error) และไม่สามารถวัดความเชื่อถือได้ (Reliability) อีกทั้งทำให้ไม่สามารถใช้ทฤษฎีทางสถิติในการอนุมานคุณลักษณะของประชากรได้ เช่น

2.1 การเลือกตัวอย่างแบบบังเอิญ (Accidental Sampling) เป็นวิธีการเลือกตัวอย่างโดยไม่มีกฎเกณฑ์ โดยจะเก็บข้อมูลจากหน่วยตัวอย่างใดก็ได้ที่เก็บได้ง่าย จนกว่าจะครบกำหนดตามที่ต้องการ ซึ่งถือความสะดวกสบายของผู้วิจัยเป็นหลัก บางครั้ง เรียกว่า **การเลือกตัวอย่างแบบสะดวก (Convenience Sampling)** เช่น เลือกเก็บข้อมูลของครัวเรือนที่อยู่ใกล้ถนนหรือทางเดิน

2.2 การเลือกตัวอย่างแบบโควตา (Quota Sampling) เป็นการเลือกตัวอย่างโดยจำแนกประชากรออกเป็นกลุ่ม ๆ ที่มีลักษณะที่สนใจที่เหมือนกัน แล้วเลือกตัวอย่างจากแต่ละกลุ่มตามขนาดตัวอย่างที่กำหนดไว้ล่วงหน้า (เรียกว่าระบบโควตา) ซึ่งไม่มีการสุ่มตัวอย่าง เพื่อให้กลุ่มตัวอย่างที่ได้กระจายอยู่ในแต่ละกลุ่ม โดยมีเงื่อนไขเพียงจำนวนรวมที่ต้องเลือกจากแต่ละกลุ่มย่อยครบตามโควตาก็ยุติการเก็บข้อมูลจากกลุ่มย่อยนั้น

2.3 การเลือกตัวอย่างตามพินิจ (Judgment Sampling) หรือ การเลือกตัวอย่างแบบเจาะจง (Purposive Sampling) เป็นการเลือกตัวอย่างที่ต้องการรายละเอียดมาก ผู้วิจัยเป็นคนตัดสินใจพิจารณาเลือกหน่วยตัวอย่าง หรือเจาะจงหน่วยตัวอย่างที่จะเก็บรวบรวมข้อมูลตามที่ต้องการ ซึ่งวิธีนี้จะใช้ประสบการณ์ และความชำนาญของผู้วิจัย เช่น การสัมภาษณ์เชิงลึก การศึกษาในเรื่องที่มีลักษณะเฉพาะกลุ่ม เช่น เลือกคนหรือลูกค้าที่คิดว่ามีลักษณะเดียวกับคนหรือลูกค้าส่วนใหญ่ การวิเคราะห์สภาพเศรษฐกิจ เป็นต้น

2.4 การเลือกตัวอย่างแบบเฉพาะหน่วยที่เลือกได้สะดวก (Accessible Sampling) เป็นการเลือกตัวอย่างโดยยึดหลักความง่าย ความสะดวกในการเลือก เช่น การเลือกผลไม้จากภาชนะที่บรรจุ การเก็บข้อมูลของผลไม้จากสวน เป็นต้น

2.5 การเลือกตัวอย่างแบบสโนว์บอลล์ (Snowball Sampling) เป็นการเลือกตัวอย่างที่ผู้วิจัยทำการสุ่มหน่วยตัวอย่างที่มีลักษณะที่สนใจ มาจำนวนหนึ่งแล้วเก็บรวบรวมข้อมูล จากนั้นให้หน่วยตัวอย่างกลุ่มนี้ แนะนำตัวอย่างที่มีความสัมพันธ์กันบางลักษณะต่อไปอีกจำนวนหนึ่ง ผู้วิจัยก็ทำการเก็บข้อมูลจากกลุ่มตัวอย่างที่ถูกแนะนำ ทำต่อไปจนได้ตัวอย่างครบตามที่ต้องการ การเลือกตัวอย่างแบบสโนว์บอลล์ จะมีการเพิ่มขนาดตัวอย่างทุกครั้ง ซึ่งเหมือนกับการกลิ้งหิมะไปรอบ ๆ จนลูกหิมะมีขนาดใหญ่ เรียกอีกอย่างหนึ่งว่า **การเลือกตัวอย่างแบบลูกโซ่**

2.6 การเลือกตัวอย่างจากประชากรที่มีการเคลื่อนไหว (Sampling of Mobile Population) เป็นการเลือกตัวอย่างจากประชากรที่มีการเคลื่อนไหว และเปลี่ยนแปลงตลอดเวลา เพื่อประมาณจำนวนประชากร เช่น ปลาบึกในแม่น้ำ สัตว์ป่า วิธีการคือ ทำการเลือกตัวอย่างมาจำนวนหนึ่งแล้วทำสัญลักษณ์ หรือเครื่องหมายแล้วปล่อยตัวอย่างคืนไปในประชากร หลังจากหน่วยตัวอย่างต่าง ๆ ที่ปล่อยไปปะปนกันดีแล้วสักระยะเวลาหนึ่งจากนั้นทำการเลือกหน่วยตัวอย่างมาจำนวนหนึ่ง และ

สังเกตว่ามีหน่วยตัวอย่างใดบ้างที่มีสัญลักษณ์ที่กำหนด วิธีการนี้เรียกอีกอย่างว่า **การเลือกตัวอย่างแบบจับ-จับใหม่ (Capture-recapture Sampling)** ตัวอย่างเช่น ต้องการประมาณจำนวนปลาบึกทั้งหมดในแม่น้ำโขง ดังนั้น จะได้ค่าประมาณปลาบึกในแม่น้ำโขง ดังนี้

$$N \approx \frac{n r}{k}$$

N = ค่าประมาณจำนวนปลาบึกในแม่น้ำโขง

n = จำนวนปลาบึกที่จับได้ในครั้งที่ 1

r = จำนวนปลาบึกที่จับได้ในครั้งที่ 2

k = จำนวนปลาบึกที่มีสัญลักษณ์ในการจับครั้งที่ 2

1.5 การทดสอบสมมติฐาน

ความหมายและชนิดของสมมติฐาน

การทดสอบสมมติฐาน (Hypothesis Testing) เป็นกระบวนการทดสอบคุณลักษณะบางประการของประชากร หรือค่าพารามิเตอร์ของประชากรว่าเป็นไปตามที่คาดหรือไม่ (การตัดสินใจเกี่ยวกับค่าของพารามิเตอร์) การกำหนดสมมติฐานนั้นมีการกำหนด 2 แบบ คือ สมมติฐานเชิงบรรยาย (Descriptive Hypothesis) ซึ่งเป็นสมมติฐานที่อยู่ในลักษณะของข้อความการบรรยาย เช่น ราคาขายของรถยนต์มือสองมีความสัมพันธ์กับอายุการใช้งานของรถยนต์ รายจ่ายของนักศึกษาหญิงมากกว่ารายจ่ายของนักศึกษาชาย และสมมติฐานเชิงสถิติ (Statistical Hypothesis) คือ ประโยคหรือข้อความเกี่ยวกับพารามิเตอร์ของหนึ่งประชากรหรือหลายประชากรที่ต้องการทดสอบว่าถูกต้องหรือไม่ โดยใช้วิธีการทางสถิติ ซึ่งนิยมเขียนอยู่ในรูปของสัญลักษณ์ทางคณิตศาสตร์ของพารามิเตอร์ เช่น อายุการใช้งานของแบตเตอรี่ในโรงงานแห่งหนึ่ง มีค่าเท่ากับ 3 ปี ให้ μ แทนอายุการใช้งานเฉลี่ยของแบตเตอรี่จากโรงงาน ดังนั้น จะตั้งสมมติฐานว่า $\mu = 3$ สมมติฐานแบ่งเป็น 2 ชนิด คือ

1) สมมติฐานจริง (สมมติฐานเพื่อการทดสอบ) หรือสมมติฐานว่าง (Null Hypothesis) แทนด้วย H_0

2) สมมติฐานแย้ง หรือสมมติฐานทางเลือก (Alternative Hypothesis) แทนด้วย H_1 หรือ H_a

การทดสอบสมมติฐานจะพิจารณาจากการแจกแจงของตัวสถิติที่ได้จากตัวอย่างในการประมาณค่าพารามิเตอร์ที่เกี่ยวข้องกับสมมติฐานนั้น ๆ ตัวอย่างเช่น ในการประมาณค่าเฉลี่ยของประชากร μ นั้น โดยปกติแล้วจะใช้ค่าเฉลี่ยของตัวอย่าง $\bar{X}$ ทำการประมาณ ทำนองเดียวกันในการทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยของประชากร μ ก็จะมีการพิจารณาจากการแจกแจงของค่าเฉลี่ยของตัวอย่าง $\bar{X}$ ในการทดสอบ ซึ่งจะเรียกว่า **ตัวสถิติทดสอบ (Test Statistic)** การที่จะยอมรับหรือปฏิเสธสมมติฐานนั้น จะพิจารณาจากค่าสถิติที่ใช้ในการทดสอบว่าตกอยู่ในอาณาเขตที่ยอมรับ H_0 หรือ ปฏิเสธ H_0 และจะเรียกอาณาเขตปฏิเสธ H_0 ว่า **เขตวิกฤต (Critical Region)** หรือ **เขตปฏิเสธ (Rejection Region)**

ซึ่งเป็นเซตของค่าสถิติทดสอบจากตัวอย่างสุ่มที่ทำให้ปฏิเสธสมมุติฐานว่าง ส่วนค่าที่กำหนดขอบเขตของอาณาเขตปฏิเสธก็จะเรียกว่า **ค่าวิกฤต (Critical Value)**

ความผิดพลาดในการทดสอบสมมุติฐาน

เนื่องจากการทดสอบสมมุติฐานจะใช้ข้อมูลจากตัวอย่างทำการอ้างอิงถึงประชากรอาจทำให้เกิดความผิดพลาดในการสรุปผล ซึ่งความผิดพลาดมี 2 แบบ

1) ความผิดพลาดแบบที่ 1 (Type I Error) คือ ความผิดพลาดที่เกิดจากการตัดสินใจปฏิเสธ H_0 โดยที่ H_0 เป็นจริง ความน่าจะเป็นที่จะเกิดความผิดพลาดชนิดนี้เรียกว่า ระดับนัยสำคัญ (Level of Significance) นั่นคือ $P(\text{ปฏิเสธ } H_0 \mid H_0 \text{ จริง}) = \alpha$

2) ความผิดพลาดแบบที่ 2 (Type II Error) คือความผิดพลาดที่เกิดจากการตัดสินใจไม่สามารถปฏิเสธ H_0 โดยที่ H_0 ไม่จริง ความน่าจะเป็นที่จะเกิดความผิดพลาดชนิดนี้แทนด้วย β นั่นคือ $P(\text{ไม่สามารถปฏิเสธ } H_0 \mid H_0 \text{ ไม่จริง}) = \beta$ และจะเรียก $(1-\beta)$ ว่า กำลังการทดสอบ (Power of a Test) ซึ่งเป็นการตัดสินใจถูกต้องที่จะปฏิเสธ H_0 โดยที่ H_0 ไม่จริง

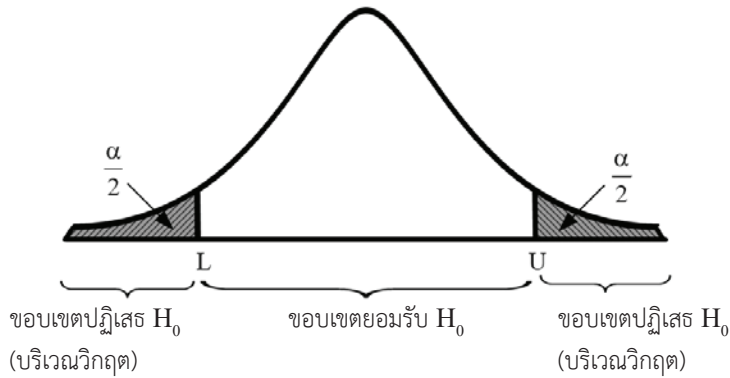
ตารางที่ 1.1 การตัดสินใจในการทดสอบสมมุติฐาน

ผลการทดสอบ	H_0 จริง	H_0 ไม่จริง
ปฏิเสธ H_0	Type I Error	ตัดสินใจถูกต้อง
ไม่สามารถปฏิเสธ H_0	ตัดสินใจถูกต้อง	Type II Error

พบว่าค่า α และ β แปรผกผันกัน คือ ถ้าค่า α มาก ค่า β จะน้อย หรือถ้าค่า α น้อย ค่า β จะมาก ในการทดสอบสมมุติฐานจะต้องพยายามควบคุมให้ค่า α และ β มีค่าน้อย ๆ วิธีการหนึ่งที่จะทำให้ α และ β มีค่าน้อยคือ การเพิ่มขนาดตัวอย่าง ในการทดสอบสมมุติฐานโดยทั่วไปแล้วจะกำหนดค่า α หรือระดับนัยสำคัญก่อนการทดสอบ

ประเภทของการทดสอบสมมุติฐาน

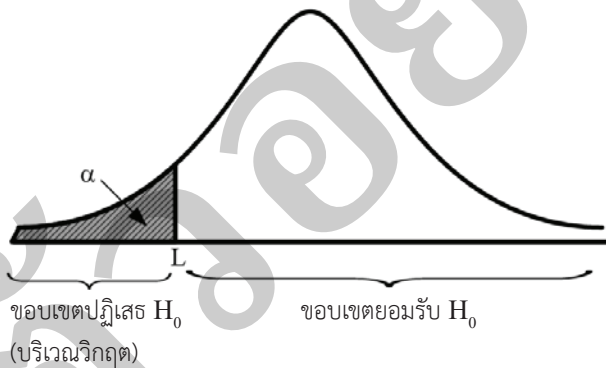
การทดสอบ 2 ด้าน (Two-sided Test) คือ สมมุติฐานที่อยู่ในรูป $H_0 : \theta = \theta_0$ เทียบกับ $H_1 : \theta \neq \theta_0$ มีขอบเขตของการยอมรับและปฏิเสธสมมุติฐานว่าง (H_0) โดย θ แทนพารามิเตอร์ที่ต้องการทดสอบสมมุติฐาน ดังรูป ซึ่งขอบเขตของการปฏิเสธ H_0 (บริเวณวิกฤต) มี 2 ด้านและค่าวิกฤตคือ L และ U



รูปที่ 1.1 ขอบเขตของการทดสอบ $H_0 : \theta = \theta_0$ เทียบกับ $H_1 : \theta \neq \theta_0$

การทดสอบด้านเดียว (One-sided Test)

ถ้าสมมติฐานอยู่ในรูป $H_0 : \theta \geq \theta_0$, $H_1 : \theta < \theta_0$ จะมีขอบเขตของการปฏิเสธ H_0 (บริเวณวิกฤต) อยู่ด้านซ้ายและค่าวิกฤตคือ L



รูปที่ 1.2 ขอบเขตของการทดสอบ $H_0 : \theta \geq \theta_0$ เทียบกับ $H_1 : \theta < \theta_0$

ถ้าสมมติฐานอยู่ในรูป $H_0 : \theta \leq \theta_0$, $H_1 : \theta > \theta_0$ จะมีขอบเขตของการปฏิเสธ H_0 (บริเวณวิกฤต) อยู่ด้านขวาและค่าวิกฤตคือ U