

สถิติและการวิจัย

ทางสังคมศาสตร์:
แนวคิด วิธีการ และการประยุกต์



ผศ.ดร.สุภัทรชัย สีสะใบ

ภาควิชารัฐศาสตร์ คณะสังคมศาสตร์ มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย

สถิติและการวิจัยทางสังคมศาสตร์:
แนวคิด วิธีการ และการประยุกต์
Statistics and Social Science Research:
Concepts, Methods and Applications

ผู้ช่วยศาสตราจารย์ ดร.สุภัทรชัย สีสะไบ
ภาควิชารัฐศาสตร์ คณะสังคมศาสตร์
มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย
พ.ศ. 2568

สถิติและการวิจัยทางสังคมศาสตร์: แนวคิด วิธีการ และการประยุกต์

Statistics and Social Science Research: Concepts, Methods and Applications

จัดทำโดย ผศ.ดร.สุภัทรชัย สีสะไบ

ภาควิชารัฐศาสตร์ คณะสังคมศาสตร์

มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย

ข้อมูลทางบรรณานุกรมของสำนักหอสมุดแห่งชาติ

Nation Library of Thailand Cataloging in Publication Data

สุภัทรชัย สีสะไบ.

สถิติวิจัยทางสังคมศาสตร์: แนวคิด วิธีการ และการประยุกต์ = Statistics and Social Science Research: Concepts, Methods and Applications.--

พระนครศรีอยุธยา: โรงพิมพ์มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย, 2568.
309 หน้า.

1. สังคมศาสตร์ --ระเบียบวิธีทางสถิติ. 2. สถิติวิเคราะห์. 3. รัฐประศาสนศาสตร์.
I. ชื่อเรื่อง

001.42

ISBN 978-616-629-365-4

ที่ปรึกษา:

พระอุดมสิทธินายก, รศ.ดร.

รศ.ดร.สุรพล สุยะพรหม

ผู้ทรงคุณวุฒิพิจารณา

รศ.ดร.เกียรติศักดิ์ สุขเหลือ้ง มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย

ผลงานวิชาการ:

รศ.ดร.ธนภุต โพธิ์เงิน มหาวิทยาลัยปทุมธานี

รศ.ดร.อัจฉรา หล่อตระกูล มหาวิทยาลัยราชภัฏพระนครศรีอยุธยา

ผู้แต่ง:

ผศ.ดร.สุภัทรชัย สีสะไบ

ออกแบบปก:

พระพลวัฒน์ สพลโ

พิมพ์ครั้งแรก:

ตุลาคม 2568

จำนวนพิมพ์:

100 เล่ม

จำนวนหน้า:

309 หน้า

ราคา:

390 บาท

พิมพ์ที่:

โรงพิมพ์มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย

79 ม.1 ต.ลำไทร อ.วังน้อย จ.พระนครศรีอยุธยา 13170

โทร. 035-001-124 ต่อ 0 มือถือ 066-060-4192

E-mail: mcupress@mcu.ac.th

คำนำ

หนังสือเล่มนี้จัดทำขึ้นเพื่อรวบรวมและนำเสนอองค์ความรู้ด้านสถิติและการวิจัยทางสังคมศาสตร์อย่างเป็นระบบ โดยมีวัตถุประสงค์เพื่อเป็นแนวทางในการศึกษาค้นคว้า วิเคราะห์ และประยุกต์ใช้หลักการทางสถิติในงานวิจัยทางสังคมศาสตร์และรัฐประศาสนศาสตร์ เนื้อหาได้เรียบเรียงตั้งแต่ความรู้พื้นฐานจนถึงการประยุกต์ใช้โปรแกรมทางสถิติ พร้อมตัวอย่างการวิจัยทางรัฐประศาสนศาสตร์และแนวปฏิบัติที่สามารถนำไปใช้ได้จริงได้

สาระสำคัญครอบคลุมทั้งด้านทฤษฎีและปฏิบัติ โดยเฉพาะการอธิบายความสัมพันธ์ระหว่างสถิติกับกระบวนการวิจัย ตลอดจนการเชื่อมโยงไปสู่การใช้ข้อมูลเชิงประจักษ์ในการวิเคราะห์และกำหนดนโยบายสาธารณะ ซึ่งถือเป็นประโยชน์อย่างยิ่งต่อแวดวงรัฐประศาสนศาสตร์และสังคมศาสตร์ในภาพรวม เนื้อหาของหนังสือเล่มนี้แบ่งออกเป็น 10 บท ดังนี้

- บทที่ 1 ความรู้เบื้องต้นเกี่ยวกับสถิติและการวิจัยทางสังคมศาสตร์
- บทที่ 2 ประเภทของข้อมูลและการเก็บรวบรวมข้อมูล
- บทที่ 3 การแจกแจงข้อมูลและสถิติเชิงพรรณนา
- บทที่ 4 ความน่าจะเป็นและการแจกแจงความน่าจะเป็น
- บทที่ 5 การอนุมานทางสถิติ
- บทที่ 6 การวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร
- บทที่ 7 การวิเคราะห์ความแตกต่างระหว่างกลุ่ม
- บทที่ 8 การใช้โปรแกรมทางสถิติในการวิเคราะห์ข้อมูล
- บทที่ 9 สถิติและการวิจัยทางรัฐประศาสนศาสตร์
- บทที่ 10 ตัวอย่างการใช้สถิติในงานวิจัยทางรัฐประศาสนศาสตร์

ผู้เขียนหวังเป็นอย่างยิ่งว่าหนังสือเล่มนี้จะเป็นประโยชน์ต่อคณาจารย์ นักวิจัย และนักศึกษา ตลอดจนผู้สนใจทั่วไป โดยสามารถใช้เป็นคู่มือและกรอบอ้างอิงเชิงวิชาการทั้งในการเรียนการสอน การทำวิทยานิพนธ์ งานวิจัยเชิงลึก และการวางแผนนโยบายสาธารณะ เพื่อส่งเสริมให้เกิดการสร้างองค์ความรู้ใหม่ การพัฒนาทักษะการวิจัยที่มีคุณภาพ และการขับเคลื่อนสังคมให้ก้าวหน้าอย่างยั่งยืนบนพื้นฐานของความรู้และคุณธรรม

ผู้ช่วยศาสตราจารย์ ดร.สุภัทรชัย สีสะไบ

สารบัญ

เรื่อง	หน้า
คำนำ	ก
สารบัญ	๒
สารบัญตาราง	ฉ
สารบัญแผนภาพ	ฎ
บทที่ 1 ความรู้เบื้องต้นเกี่ยวกับสถิติและการวิจัยทางสังคมศาสตร์	1
1. ความหมายและประเภทของสถิติ	1
2. ประเภทของงานวิจัยทางสังคมศาสตร์	3
3. บทบาทของสถิติในการวิจัยทางสังคมศาสตร์	6
4. กระบวนการวิจัยและตำแหน่งของสถิติ	8
5. ธรรมชาติของข้อมูลทางสังคม	11
6. จริยธรรมในการใช้สถิติ	14
สรุปท้ายบท	17
บทที่ 2 ประเภทของข้อมูลและการเก็บรวบรวมข้อมูล	19
1. ประเภทของข้อมูล	20
2. ระดับของข้อมูล	22
3. การเก็บและรวบรวมข้อมูล	27
4. การตรวจสอบคุณภาพของเครื่องมือที่ใช้ในการวิจัย	38
5. การสุ่มตัวอย่างและการกำหนดขนาดตัวอย่าง	44
สรุปท้ายบท	52
บทที่ 3 การแจกแจงข้อมูลและสถิติเชิงพรรณนา	54
1. ตารางแจกแจงความถี่ (Frequency Distribution)	54
2. ค่ากลาง (Measures of Central Tendency)	57
3. ค่าการกระจาย (Measures of Dispersion)	62
4. การวิเคราะห์รูปร่างการแจกแจง (Skewness & Kurtosis)	64
5. การจำแนกข้อมูลและลักษณะการแจกแจง	73
6. การนำเสนอข้อมูลด้วยกราฟ	78
สรุปท้ายบท	91

สารบัญ (ต่อ)

เรื่อง		หน้า
บทที่ 4	ความน่าจะเป็นและการแจกแจงความน่าจะเป็น	92
	1. ความหมายและแนวคิดของความน่าจะเป็น	93
	2. กฎของความน่าจะเป็น	96
	3. การแจกแจงแบบไม่ต่อเนื่อง (Discrete Distributions)	100
	4. การแจกแจงแบบต่อเนื่อง (Continuous Distribution)	102
	5. การใช้ตาราง Z	106
	6. ความสำคัญของความน่าจะเป็นต่อการวิเคราะห์เชิงอนุมาน	108
	สรุปท้ายบท	109
บทที่ 5	การอนุมานทางสถิติ	111
	1. ประชากรและตัวอย่าง (Population vs Sample)	112
	2. การประมาณค่าเฉลี่ยและสัดส่วนของประชากร	117
	3. ช่วงความเชื่อมั่น (Confidence Interval)	120
	4. การทดสอบสมมติฐาน (Hypothesis Testing)	122
	5. ค่าพี (p-value) และระดับนัยสำคัญ	124
	6. การทดสอบค่าเฉลี่ย (Mean Testing)	127
	7. ความผิดพลาดแบบที่ 1 และแบบที่ 2 (Type I & Type II Errors)	129
	สรุปท้ายบท	132
บทที่ 6	การวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร	133
	1. ความหมายของความสัมพันธ์ (Correlation/ Relationship)	133
	2. สหสัมพันธ์เพียร์สัน (Pearson's Product-Moment Correlation, r)	136
	3. สหสัมพันธ์แบบสเปียร์แมน (Spearman's Rank Correlation Coefficient, p)	138
	4. การถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression)	140
	5. การพล็อตกราฟแนวโน้ม (Trend Plotting)	146
	6. ข้อจำกัดและสมมติฐานของการวิเคราะห์	149
	คำถามท้ายบท	151

สารบัญ (ต่อ)

เรื่อง		หน้า
บทที่ 7	การวิเคราะห์ความแตกต่างระหว่างกลุ่ม	153
	1. การเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม	153
	2. t-test สำหรับสองกลุ่มอิสระ (Independent Samples t-test)	155
	3. t-test สำหรับกลุ่มที่วัดซ้ำ (Paired Samples t-test)	158
	4. การวิเคราะห์ความแปรปรวน (Analysis of Variance: ANOVA)	161
	5. ข้อกำหนดเบื้องต้นและการตรวจสอบสมมติฐาน	168
	สรุปท้ายบท	170
บทที่ 8	การใช้โปรแกรมทางสถิติในการวิเคราะห์ข้อมูล	171
	1. แนะนำโปรแกรม SPSS, JASP, R และ PSPP	171
	2. การป้อนข้อมูลและการจัดโครงสร้าง Dataset	178
	3. การวิเคราะห์ข้อมูลเชิงพรรณนา (Descriptive Statistics)	182
	4. การวิเคราะห์ข้อมูลเชิงอนุมาน (Inferential Statistics)	184
	5. การแปลผล Output	202
	6. การนำเสนอผลลัพธ์ในรูปแบบข้อความ ตาราง และกราฟ	204
	7. การเขียนรายงานผลการวิเคราะห์	208
	8. ข้อควรระวังในการใช้โปรแกรมทางสถิติ	216
	สรุปท้ายบท	218
บทที่ 9	สถิติและการวิจัยทางรัฐประศาสนศาสตร์	219
	1. ความหมายและความสำคัญของสถิติและการวิจัยทางรัฐประศาสนศาสตร์	220
	2. ประเภทของสถิติและวิธีการวิจัยทางรัฐประศาสนศาสตร์	221
	3. กระบวนการวิจัยทางรัฐประศาสนศาสตร์	223
	4. การประยุกต์ใช้สถิติในการวิจัยและการประเมินนโยบายสาธารณะ	225
	5. กรณีศึกษาและตัวอย่างการใช้สถิติและการวิจัยทางรัฐประศาสนศาสตร์	227
	6. ข้อจำกัดของสถิติและการวิจัยทางรัฐประศาสนศาสตร์	228
	7. จริยธรรมในการวิจัยทางรัฐประศาสนศาสตร์	230
	8. เครื่องมือและเทคโนโลยีสมัยใหม่ในการวิจัยทางรัฐประศาสนศาสตร์	231

สารบัญ (ต่อ)

เรื่อง	หน้า
9. แนวโน้มและทิศทางอนาคตของการวิจัยและการใช้สถิติทาง รัฐประศาสนศาสตร์	233
สรุปท้ายบท	234
บทที่ 10 ตัวอย่างการใช้สถิติในงานวิจัยทางรัฐประศาสนศาสตร์	236
1. การใช้สถิติเชิงพรรณนา (Descriptive Statistics)	237
2. การใช้สถิติเชิงอนุมาน (Inferential Statistics)	242
3. การใช้สถิติขั้นสูงในงานวิจัยรัฐประศาสนศาสตร์	257
สรุปท้ายบท	285
บรรณานุกรม	287
ประวัติผู้แต่ง	295

สารบัญตาราง

ตารางที่	หน้า
2.1	สรุประดับของข้อมูล (Scales of Measurement)
2.2	เปรียบเทียบการเก็บรวบรวมข้อมูลเชิงคุณภาพ
2.3	เปรียบเทียบการเก็บรวบรวมข้อมูลเชิงปริมาณ
2.4	รูปแบบของแบบตรวจสอบที่ให้ผู้เชี่ยวชาญพิจารณาความเที่ยงตรงเชิงเนื้อหาของเครื่องมือที่ใช้ในการวิจัย
2.5	ผลการวิเคราะห์ IOC
2.6	ผลการทดสอบความเชื่อถือ
2.7	การสุ่มที่ใช้ความน่าจะเป็น
2.8	การสุ่มที่ไม่ใช้ความน่าจะเป็น
2.9	ขนาดของกลุ่มตัวอย่างของเครชีและมอร์แกนที่ระดับความเชื่อมั่น 95%
3.1	ตารางแจกแจงความถี่ (Ungrouped)
3.2	คะแนนสอบนักเรียน
3.3	ข้อมูลคะแนนสอบนักเรียน
3.4	สรุปเปรียบเทียบประเภท Kurtosis
3.5	สรุปแนวทางเลือกใช้กราฟ
3.6	ข้อมูลสรุป (ช่วงคะแนนแบบ Bar/Pie)
3.7	Histogram (bins = 6 จาก min-max)
5.1	การเลือกใช้ Z-test หรือ t-test
6.1	ชั่วโมงอ่านหนังสือและคะแนนสอบของนักเรียน
6.2	ความถี่การใช้ห้องสมุดและผลการประเมินความพึงพอใจของนักศึกษา
6.3	ชั่วโมงอ่านหนังสือและคะแนนสอบของนักเรียน
6.4	ตัวอย่างการพล็อตกราฟชั่วโมงอ่านหนังสือและคะแนนสอบของนักเรียน
7.1	คะแนนทดสอบก่อนเรียนและหลังเรียน
7.2	ผลสัมฤทธิ์ทางการเรียนของนักเรียน
7.3	ผลการเปรียบเทียบผลสัมฤทธิ์ทางการเรียน
7.4	ค่าเฉลี่ยผลสัมฤทธิ์ทางการเรียนของนักเรียน
8.1	Variable View
8.2	Codebook

สารบัญตาราง (ต่อ)

ตารางที่	หน้า
8.3 Output SPSS ค่ากลาง	182
8.4 Output SPSS ความถี่และร้อยละ	182
8.5 Descriptive Statistics	184
8.6 One-Sample Statistics	186
8.7 One-Sample Test	186
8.8 Group Statistics	187
8.9 Independent Samples Test	187
8.10 คะแนนสอบของนักเรียน	188
8.11 Paired Samples Statistics	188
8.12 Paired Samples Test	189
8.13 ระดับการศึกษา (ต่ำ-กลาง-สูง)	189
8.14 ANOVA	190
8.15 Post Hoc Tests (Tukey HSD)	190
8.16 เพศกับวิธีการสอน	191
8.17 Tests of Between-Subjects Effects	192
8.18 จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์และคะแนนสอบปลายภาคของ นักศึกษา	193
8.19 Tests of Between-Subjects Effects	194
8.20 ระดับความพึงพอใจต่อการเรียน	195
8.21 Tests of Between-Subjects Effects	195
8.22 จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์	196
8.23 Model Summary	197
8.24 ANOVA Table	197
8.25 Coefficients Table	197
8.26 จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์	198
8.27 Model Summary	199
8.28 ANOVA Table	199
8.29 Coefficients Table	199

สารบัญตาราง (ต่อ)

ตารางที่	หน้า
8.30 Crosstab (Observed frequencies)	200
8.31 Chi-Square Tests	200
8.32 ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของคะแนนก่อน-หลังเรียน	204
8.33 Group Statistics	209
8.34 Independent Samples Test (Welch's t-test)	210
8.35 Paired Samples Statistics	210
8.36 Paired Samples Test	210
8.37 One-sample Statistics	210
8.38 One-sample Test	210
8.39 Descriptive Statistics (ตามระดับการศึกษา)	211
8.40 ANOVA Summary	211
8.41 Post hoc (Tukey HSD)	211
8.42 Tests of Between-Subjects Effects	212
8.43 Pearson Correlation Matrix	212
8.44 Spearman's Correlation Matrix	212
8.45 Model Summary	213
8.46 Coefficients (Simple Regression)	213
8.47 Model Summary	213
8.48 Coefficients (Multiple Regression)	213
8.49 Crosstab เพศ × การเลือกวิชาเรียน	214
8.50 Chi-square Tests	214
10.1 จำนวนและร้อยละของผู้ตอบแบบสอบถามจำแนกตามปัจจัยส่วนบุคคล	238
10.2 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และระดับการปฏิบัติตามการบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ โดยภาพรวม	241
10.3 การเปรียบเทียบระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำ ตาเสา อำเภอลำดวน จังหวัดพระนครศรีอยุธยา ภาพรวม จำแนกตาม เพศ	243

สารบัญตาราง (ต่อ)

ตารางที่	หน้า
10.4 การเปรียบเทียบระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำ ตาเสา อำเภองาวน้อย จังหวัดพระนครศรีอยุธยา ภาพรวม จำแนกตาม ระดับการศึกษา	244
10.5 การเปรียบเทียบความแตกต่างรายคู่ด้วยวิธีผลต่างนัยสำคัญน้อยที่สุด ของการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภองาว น้อย จังหวัดพระนครศรีอยุธยา ด้านการพัฒนา จำแนกตามระดับ การศึกษา	245
10.6 ความสัมพันธ์ระหว่างปัจจัยเชิงใจหรือปัจจัยที่เป็นตัวกระตุ้นในการ ปฏิบัติงาน ปัจจัยค้ำจุนหรือปัจจัยสุขอนามัย และหลักอิทธิบาท 4 กับ ประสิทธิภาพในการปฏิบัติงานของบุคลากรมหาวิทยาลัยมหาจุฬาลง กรณราชวิทยาลัย	247
10.7 การทดสอบสัมประสิทธิ์สหสัมพันธ์ระหว่างการglomเกลตาทางการเมือง และหลักสังคหวัตถุ 4	251
10.8 ค่า Tolerance และค่า Variance Inflation Factor (VIF) ของการ glomเกลตาทางการเมืองและหลักสังคหวัตถุ 4	252
10.9 การglomเกลตาทางการเมือง ประกอบด้วย 1) ด้านครอบครัว 2) ด้าน กลุ่มเพื่อน 3) ด้านสถาบันการศึกษา 4) ด้านกลุ่มทางสังคม อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรค การเมืองในจังหวัดพระนครศรีอยุธยา	253
10.10 หลักสังคหวัตถุ 4 ประกอบด้วย 1) ทาน โอบอ้อมอารี 2) ปิยวาจา วจีไพเราะ 3) อุตถจริยา สงเคราะห์ประชาชน 4) สมานัตตตา วางตนพอดี อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรค การเมืองในจังหวัด พระนครศรีอยุธยา	255
10.11 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน ระหว่างตัวแปรในองค์ประกอบหลักการบริหารจัดการแบบมีส่วนร่วม	260
10.12 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดองค์ประกอบ หลักการบริหารจัดการแบบมีส่วนร่วม	261

สารบัญตาราง (ต่อ)

ตารางที่	หน้า
10.13 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สันระหว่างตัวแปรในองค์ประกอบปัจจัยแห่งความสำเร็จ	263
10.14 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดปัจจัยแห่งความสำเร็จ	264
10.15 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สันระหว่างตัวแปรในองค์ประกอบธรรมมะเพื่อการปฏิบัติ	265
10.16 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดองค์ประกอบธรรมมะเพื่อการปฏิบัติ	266
10.17 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สันระหว่างตัวแปรในองค์ประกอบประสิทธิผลพุทธิวิธีการบริหารจัดการศูนย์ฯ	267
10.18 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดองค์ประกอบประสิทธิผลพุทธิวิธีการบริหารจัดการ	268
10.19 ค่าสถิติเบื้องต้นของตัวแปรสังเกตได้ในโมเดลประสิทธิผลพุทธิวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร	274
10.20 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และค่าสัมประสิทธิ์สหสัมพันธ์แบบเพียร์สันของตัวแปรสังเกตได้ในโมเดลประสิทธิผลพุทธิวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ที่มีตัวแปรธรรมมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน	278
10.21 ค่าสถิติผลการวิเคราะห์แยกค่าสหสัมพันธ์ระหว่างตัวแปรแฝงและการวิเคราะห์หือทธิพลของโมเดลประสิทธิผลพุทธิวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตรที่มีตัวแปรธรรมมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน	283

สารบัญแผนภาพ

แผนภาพที่	หน้า
2.1 ความสัมพันธ์ของระดับข้อมูล	26
3.1 Positive Skew	65
3.2 Negative Skew	66
3.3 Symmetrical	66
3.4 Histogram with Bell-shaped Curve	68
3.5 Mesokurtic	69
3.6 Leptokurtic	69
3.7 Platykurtic	70
3.8 Histogram of Example Scores (Leptokurtic: Kurtosis > 0)	71
3.9 Histogram of Example Scores (Platykurtic: Kurtosis < 0)	72
3.10 Symmetrical Distribution	74
3.11 Positive Skew	75
3.12 Negative Skew	76
3.13 Bar Chart (Gender Distribution)	79
3.14 Pie Chart (Proportions)	80
3.15 Histogram (Data Distribution)	81
3.16 Line Graph (Monthly Income)	82
3.17 Bar Chart	88
3.18 Pie Chart	89
3.19 Histogram	90
3.20 Line Graph	90
4.1 Venn Diagram	97
4.2 Tree Diagram	98
4.3 การแจกแจงทวินาม (Binomial Distribution)	100
4.4 การแจกแจงปัวส์ซอง (Poisson Distribution)	101
4.5 การแจกแจงปกติ (Normal Distribution)	103
4.6 แสดงการใช้ตาราง Z	107
5.1 การแจกแจง t-distribution	124

สารบัญแผนภาพ (ต่อ)

แผนภาพที่	หน้า
5.2 Type I Error	130
5.3 Type II Error	131
6.1 Scatter Plot + Regression Line	147
8.1 โปรแกรม SPSS	173
8.2 โปรแกรม JASP	175
8.3 โปรแกรม RStudio	176
8.4 โปรแกรม PSPP	178
8.5 Histogram	183
8.6 Bar chart	183
8.7 Pie chart	184
8.8 Interaction Plot	192
8.9 ตัวอย่างผลการทดสอบไคสแควร์ (Chi-square Test)	201
8.10 แผนภาพตัวอย่าง Histogram	205
8.11 แผนภาพตัวอย่าง Bar Chart	206
8.12 แผนภาพตัวอย่าง Boxplot	206
8.13 แผนภาพตัวอย่าง Scatter Plot	207
8.14 แผนภาพตัวอย่าง Line Chart	207
10.1 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยหลักการบริหารจัดการแบบมีส่วนร่วม	262
10.2 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยแห่งความสำเร็จ	264
10.3 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยธรรมะเพื่อการปฏิบัติ	266
10.4 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์	268
10.5 โมเดลประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตรที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน (ปรับโมเดลแล้ว)	284

บทที่ 1

ความรู้เบื้องต้นเกี่ยวกับสถิติและการวิจัยทางสังคมศาสตร์ (INTRODUCTION TO STATISTICS AND SOCIAL SCIENCE RESEARCH)

สถิติและการวิจัยทางสังคมศาสตร์นับเป็นเครื่องมือสำคัญที่ช่วยให้นักวิชาการและนักวิจัยสามารถเข้าใจปรากฏการณ์ต่าง ๆ ของมนุษย์และสังคมได้อย่างมีระบบและเป็นเหตุเป็นผล ในโลกยุคปัจจุบันที่เต็มไปด้วยข้อมูลจำนวนมาก การใช้สถิติทำให้ข้อมูลที่กระจัดกระจายสามารถถูกรวบรวม วิเคราะห์ และตีความออกมาอย่างเป็นระบบ เพื่อนำไปสู่ข้อสรุปที่มีความน่าเชื่อถือและสามารถใช้ประกอบการตัดสินใจได้อย่างมีประสิทธิภาพ ขณะเดียวกัน การวิจัยทางสังคมศาสตร์ก็เป็นวิธีการค้นหาความจริงเกี่ยวกับพฤติกรรม ความคิด ความสัมพันธ์ และปฏิสัมพันธ์ของมนุษย์ในสังคม ซึ่งเป็นฐานสำคัญต่อการแก้ไขปัญหาสังคม การพัฒนานโยบาย และการสร้างองค์ความรู้ใหม่

การทำความเข้าใจเรื่องสถิติและการวิจัยในสังคมศาสตร์จึงมีความสำคัญยิ่ง ทั้งในเชิงทฤษฎีและการปฏิบัติ ไม่ว่าจะเป็นการรู้จักความหมายของสถิติ ประเภทของสถิติทั้งเชิงพรรณนาและเชิงอนุมาน ประเภทของการวิจัยเชิงคุณภาพและเชิงปริมาณ รวมถึงบทบาทของสถิติในกระบวนการวิจัย ตั้งแต่การออกแบบ เก็บรวบรวมข้อมูล วิเคราะห์ ไปจนถึงการนำเสนอและแปลผลข้อมูล นอกจากนี้ การพิจารณาถึงธรรมชาติของข้อมูลทางสังคมที่มีความซับซ้อนและหลากหลาย ตลอดจนการตระหนักถึงจริยธรรมในการใช้สถิติและการทำวิจัย ล้วนเป็นหัวใจสำคัญที่จะทำให้นักวิจัยมีคุณค่าและน่าเชื่อถือ

ดังนั้น บทนี้จึงมุ่งอธิบายให้ผู้อ่านเห็นภาพรวมของ “สถิติและการวิจัยทางสังคมศาสตร์” ตั้งแต่ความหมาย ประเภท บทบาท และกระบวนการ ไปจนถึงประเด็นด้านคุณธรรมและจริยธรรม เพื่อเป็นพื้นฐานในการเรียนรู้และเป็นแนวทางสำหรับการศึกษาและการทำวิจัยในสาขาสังคมศาสตร์ต่อไป

1. ความหมายและประเภทของสถิติ

คำว่า สถิติ ตรงกับภาษาอังกฤษว่า Statistics ซึ่งมีรากศัพท์มาจากคำว่า State ความหมายเดิมจึงหมายถึง ข้อมูล (data) หรือข่าวสาร (information) ที่เป็นประโยชน์แก่รัฐหรือประเทศในด้านต่าง ๆ เช่น ข้อมูลในการบริหารงานหรือวางแผนเกี่ยวกับกำลังคน การเก็บภาษีอากรเพื่อเป็นรายได้ของรัฐ การเกณฑ์ ทหารเพื่อเข้าประจำการรักษาความปลอดภัยและป้องกันประเทศ การจัดการศึกษา การประกันสังคม และการสาธารณสุข

เป็นต้น ต่อมาความหมายของคำว่า สถิติหมายถึงรวมถึงการค้นคว้าและพัฒนาในด้านเนื้อหา และวิธีการของนักคณิตศาสตร์และนักสถิติจำนวนมาก จึงอาจสรุปความหมายของสถิติ หมายถึง ตัวเลขหรือข้อความจริงต่าง ๆ ที่จัดบันทึกไว้เป็นหลักฐานอาจเป็นตัวเลขที่ใช้ บรรยายเหตุการณ์หรือข้อเท็จจริงของเรื่องต่าง ๆ ที่เราต้องการศึกษา เช่น สถิติจำนวน ผู้ป่วย สถิติจำนวนคนเกิด สถิติจำนวนคนตาย เป็นต้น หรือกระบวนการที่เกี่ยวข้องกับ ตัวเลขและข้อความจริงที่ถูกสรุปและตีความโดยกระบวนการทางสถิติ (กรมวิชาการเกษตร, 2558)

จากความหมายเหล่านี้จึงกล่าวโดยสรุปว่า สถิติ (Statistics) คือสาขาหนึ่งของ คณิตศาสตร์ประยุกต์ที่เกี่ยวข้องกับการรวบรวม (Collection) การบรรยาย (Description) การวิเคราะห์ (Analysis) และการตีความข้อมูล (Interpretation of data) ที่ได้มาจาก กลุ่มตัวอย่างของประชากรขนาดใหญ่ การสุ่มตัวอย่างทางสถิติถูกนำมาใช้ในหลากหลาย สาขา เช่น การแพทย์ การเงิน การตลาด และอีกมากมาย เพื่อเพิ่มความเข้าใจและ ประกอบการตัดสินใจ (Chappelow, J., 2025)

โดยแบ่งความหมายของคำว่า สถิติ ได้ 3 ความหมาย ดังนี้

1. สถิติ คือ ตัวเลขที่ใช้บรรยายเหตุการณ์หรือข้อเท็จจริง (facts) ของเรื่องต่าง ๆ ที่เราต้องการศึกษา เช่น สถิติจำนวนผู้ป่วย สถิติจำนวนคนเกิด สถิติจำนวนคนตาย เป็นต้น
2. สถิติ คือ ศาสตร์หรือวิชาที่ว่าด้วยหลักการและระเบียบวิธีทางสถิติ สถิติใน ความหมายนี้มักเรียกว่า สถิติศาสตร์ (Statistics)
3. สถิติ คือ ค่าที่คำนวณขึ้นมาจากตัวอย่าง เพื่อแสดงถึงคุณลักษณะบางอย่าง ของข้อมูลชุดนั้น โดยทั่วไปจะนำค่าสถิติไปใช้ในการประมาณ ค่าพารามิเตอร์ ตัวอย่างเช่น ถ้าเราสนใจรายได้เฉลี่ยของคนในหมู่บ้านแล้ว เราสามารถนำรายได้ของทุกคนมารวมกัน แล้วหาค่าเฉลี่ยของรายได้ ค่าเฉลี่ยที่คำนวณได้นี้ถือว่าเป็นค่าพารามิเตอร์ แต่ถ้าเราสุ่ม ตัวอย่างคนในหมู่บ้านมาจำนวนหนึ่งแล้วคำนวณรายได้เฉลี่ยค่าเฉลี่ยที่ได้นี้จะเป็นค่าสถิติ (วรวรรณ วงศ์บุตร, 2564)

ดังนั้น สถิติแบ่งออกเป็น 2 ประเภทหลัก ๆ ได้แก่

1. สถิติเชิงพรรณนา (Descriptive Statistics) วิธีการทางสถิติที่ใช้ในการ สรุป บรรยาย อธิบายลักษณะข้อมูล que ศึกษาโดยรวบรวมจากกลุ่มตัวอย่างหรือประชากร ทั้งหมด เพื่อให้เข้าใจภาพรวมของข้อมูลได้ง่ายขึ้น ไม่สามารถคาดคะเนลักษณะต่าง ๆ นอกเหนือไปจากข้อมูลที่มีอยู่ได้ โดยการนำเสนอข้อมูลมักจะอยู่ในรูปแบบ เช่น

- 1) การวัดแนวโน้มเข้าสู่ส่วนกลาง (Measures of Central Tendency) เช่น ค่าเฉลี่ย (Mean) มัธยฐาน (Median) และฐานนิยม (Mode)

2) การวัดการกระจายของข้อมูล (Measures of Dispersion) เช่น ค่าพิสัย (Range) ส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation) และความแปรปรวน (Variance)

3) การนำเสนอในรูปแบบกราฟและตาราง เช่น กราฟแท่ง (Bar Chart) แผนภูมิวงกลม (Pie Chart) และตารางแจกแจงความถี่ (Frequency Distribution Table)

2. สถิติเชิงอนุมาน (Inferential Statistics) วิธีการทางสถิติที่มีความสำคัญและถูกใช้มากกว่าสถิติเชิงพรรณนา ใช้อธิบายคุณลักษณะของสิ่งที่ต้องการศึกษากลุ่มใดกลุ่มหนึ่งหรือหลายกลุ่มแล้วสามารถอ้างอิงไปยังกลุ่มประชากรได้ แต่กลุ่มตัวอย่างที่นำมาศึกษาจะต้องเป็นตัวแทนที่ดีของประชากรและได้มาโดยวิธีการสุ่มตัวอย่าง ซึ่งสถิติอ้างอิงสามารถแบ่งออกเป็น 2 ประเภท คือ (กรมวิชาการเกษตร, 2558)

1) สถิติมีพารามิเตอร์ (Parametric Statistics) เป็นวิธีการทางสถิติที่จะต้องเป็นไปตามข้อตกลงเบื้องต้นต้องมีคุณลักษณะ ดังนี้

- (1) ระดับการวัดของข้อมูลควรจะเป็นระดับช่วง (Interval) หรืออัตราส่วน (Ratio)
- (2) ข้อมูลต้องมีการแจกแจงเป็นโค้งปกติ (Normal distribution)
- (3) กลุ่มประชากรแต่ละกลุ่มที่นำมาศึกษาต้องมีความแปรปรวนเท่ากัน และตัวอย่างที่ถูกเลือกมาควรจะต้องมีความเป็นอิสระต่อกัน สถิติประเภทนี้ เช่น t-test, Z-test, ANOVA, Regression ฯลฯ

2) สถิติไร้พารามิเตอร์ (Nonparametric Statistics) เป็นวิธีการทางสถิติที่สามารถนำมาใช้ได้โดยปราศจากข้อตกลงเบื้องต้น สถิติที่อยู่ในประเภทนี้ เช่น Chi-square test, Median test, Sign test ฯลฯ

โดยปกติแล้วนักวิจัยนิยมใช้สถิติมีพารามิเตอร์ ทั้งนี้เพราะผลลัพธ์ที่ได้มีอำนาจทดสอบ (Power of Test) สูงกว่าการใช้สถิติไร้พารามิเตอร์ ดังนั้นหากข้อมูลมีคุณสมบัติที่สอดคล้องกับข้อตกลงเบื้องต้นให้ใช้สถิติมีพารามิเตอร์

2. ประเภทของงานวิจัยทางสังคมศาสตร์

การวิจัยทางสังคมศาสตร์ หมายถึง การศึกษาค้นคว้าหาความจริงด้วยระบบและวิธีการทาง วิทยาศาสตร์เกี่ยวกับพฤติกรรม ปรากฏการณ์ หรือปฏิบัติการ ตลอดจนความรู้สึกนึกคิดของมนุษย์และสังคม เพื่อให้ทราบถึงความรู้และความจริงที่จะนำมาแก้ไข ปัญหาของสังคม หรือก่อให้เกิดความรู้ใหม่ (สำนักงานคณะกรรมการวิจัยแห่งชาติ, 2561) โดยแนวทางการวิจัยทางสังคมศาสตร์ มีเทคนิคในการรวบรวม และวิเคราะห์ข้อมูล ซึ่งได้รับอิทธิพลจากกระบวนการทัศน์ (วิธีคิด วิธีมองโลกและปรากฏการณ์ รวมถึงเรื่องของความเชื่อ และโลกทัศน์ที่ตกผลึกอยู่ในวิธีคิดของบุคคลหรือของวัฒนธรรม มีอิทธิพลต่อ

พฤติกรรมและการแก้ปัญหาในชีวิตของบุคคลและสังคม) และแนวคิดทฤษฎี ซึ่งการวิจัยเชิงปริมาณและเชิงคุณภาพจะมีเทคนิค กระบวนทัศน์ และแนวคิดทฤษฎีเฉพาะแตกต่างกัน (พิมลพรรณ อิศรภักดี, 2557)

การวิจัยทางสังคมศาสตร์ สามารถแบ่งตามลักษณะการวิเคราะห์ข้อมูล แบ่งเป็น 2 ประเภท ดังนี้

1. การวิจัยเชิงคุณภาพ (Qualitative research) เป็นการวิจัยที่นำเอาข้อมูลทางด้านคุณภาพมาวิเคราะห์ ข้อมูลทางด้านคุณภาพเป็นข้อมูลที่ไม่เป็นตัวเลขแต่จะเป็นข้อความบรรยายลักษณะสภาพเหตุการณ์ของสิ่งต่าง ๆ ที่เกี่ยวข้อง และการเสนอผลการวิจัยก็จะออกมาในรูปของข้อความที่ไม่มีตัวเลขทางสถิติสนับสนุนเช่นเดียวกัน การวิจัยประเภทนี้จึงมุ่งบรรยายหรืออธิบายเหตุการณ์ต่าง ๆ โดยอาศัยความคิดวิเคราะห์ เพื่อประเมินผลหรือสรุปผลนั่นเอง

2. การวิจัยเชิงปริมาณ (Quantitative research) เป็นการวิจัยที่นำเอาข้อมูลเชิงปริมาณมาวิเคราะห์ กล่าวคือใช้ตัวเลขประกอบการวิเคราะห์ สรุปผล และการเสนอผลการวิจัยก็ออกมาเป็นตัวเลขเช่นเดียวกัน ดังนั้น การวิจัยประเภทนี้จึงมุ่งที่จะอธิบายเหตุการณ์ต่าง ๆ โดยอาศัยตัวเลขยืนยันยืนยันแสดงปริมาณมากน้อยแทนที่จะใช้ข้อความบรรยายให้เหตุผล (จิรภัทร เตชะกุลชัยนันต์, 2555)

อนึ่ง สามารถจำแนกความแตกต่างระหว่างการวิจัยเชิงปริมาณและการวิจัยเชิงคุณภาพ ได้ดังนี้

1. การวิจัยเชิงปริมาณ เป็นเรื่องของการหาความสัมพันธ์ระหว่างตัวแปร อย่างน้อย 2 ตัว เพื่อตรวจสอบสมมติฐานที่ได้เจาะจงตั้งเอาไว้ก่อน โดยรองรับด้วยแนวคิดทฤษฎีหรือองค์ความรู้ต่าง ๆ แต่การวิจัยเชิงคุณภาพจะเป็นเรื่องปรากฏการณ์ทางสังคม เพื่อทำความเข้าใจและอธิบายความหมายปรากฏการณ์ทางสังคม ซึ่งต้องดูเป็นองค์รวม (holistic) เพราะชีวิตคนหรือสังคมมีเรื่องที่เกี่ยวข้องกันหลายเรื่องไม่สามารถดูตัวแปร 2 - 3 ตัวได้ การวิจัยเชิงคุณภาพจึงไม่จำเป็นต้องตั้งสมมติฐานหรือมีแนวคิดทฤษฎีรองรับเอาไว้ก่อน แต่เป็นการสร้างองค์ความรู้หรือทฤษฎีใหม่ ๆ ตลอดจนข้อเท็จจริงใหม่จากที่เคยรู้มาแต่เดิม

2. การวิจัยเชิงปริมาณ ไม่ให้ความสนใจในบริบทรอบ ๆ ว่าเป็นอย่างไร เพราะสามารถควบคุมตัวแปรได้หมด แต่การวิจัยเชิงคุณภาพให้ความสนใจในเรื่องของบริบท (context) ทางสังคมวัฒนธรรม เพราะบริบทในแต่ละแห่งไม่เหมือนกัน เช่น ในเมืองกับในชนบท พุทธกับมุสลิม เป็นต้น

3. การวิจัยเชิงปริมาณ เก็บรวบรวมข้อมูลโดยใช้แบบสอบถามเป็นหลัก ใช้ระยะเวลาศึกษาไม่นานไม่ต้องทำความรู้จักหรือสร้างความคุ้นเคยสนิทสนมก่อน เมื่อตอบ

แบบสอบถามให้เสร็จเรียบร้อยแล้วผู้วิจัยก็จากไปหรือที่เรียกว่า “ตีหัวแล้ววิ่งหนี” ในขณะที่แบบสอบถามอาจมีข้อจำกัด เช่น การเก็บรวบรวมข้อมูลกับผู้ที่ไม่รู้หนังสือหรือมีการศึกษาน้อยเขาจะไม่สามารถตอบแบบสอบถามได้ หรือถ้าแบบสอบถามมีจำนวนมากคนที่ก็ไม่อยากตอบแบบสอบถาม แต่การวิจัยเชิงคุณภาพเก็บรวบรวมข้อมูลโดยการศึกษาเอกสาร การสังเกต การสัมภาษณ์ การสัมภาษณ์แบบเจาะลึก การตะล่อมกล่อมเกล่า การสนทนากลุ่ม ซึ่งผู้วิจัยต้องออกไปสัมผัสแหล่งข้อมูลด้วยตนเองกับกลุ่มคนที่เป็นเจ้าของปัญหา ในเบื้องต้นจะต้องสร้างความคุ้นเคยสนิทสนมก่อนและใช้การศึกษาติดตามระยะยาว

4. การวิจัยเชิงปริมาณ ลักษณะข้อมูลที่ได้จะเป็นตัวเลขหรือสถิติ สามารถเจนนับได้ แต่การวิจัยเชิงคุณภาพ ลักษณะข้อมูลจะเป็นการพรรณนาความเกี่ยวกับประวัติความเป็นมา สภาพแวดล้อม บริบททางสังคมวัฒนธรรม รากเหง้าของปัญหา ความรู้สึกนึกคิด การให้ความหมายหรือคุณค่ากับสิ่งต่าง ๆ ตลอดจนค่านิยมพฤติกรรมหรืออุดมการณ์ของบุคคล

5. การวิจัยเชิงปริมาณ ตรวจสอบเครื่องมือในการวิจัย คือ แบบสอบถามหรือแบบทดสอบ โดยผู้เชี่ยวชาญที่เป็นนักวิชาการ เพื่อหาความเที่ยงตรงเชิงเนื้อหา (content validity) และนำแบบสอบถามไปทดลองใช้ (try out) กับกลุ่มตัวอย่างที่ใกล้เคียงกันเพื่อหาค่าความเชื่อมั่น (reliability) ของเครื่องมือที่ใช้ในการวัด แต่การวิจัยเชิงคุณภาพไม่จำเป็นต้องตรวจสอบเครื่องมือในการวิจัย เพราะตัวผู้วิจัยถือเป็นเครื่องมือที่สำคัญ ทำการตรวจสอบข้อมูลที่ได้มาโดยวิธีการที่เรียกว่า การตรวจสอบแบบสามเส้า (triangulation) ได้แก่ 1) การตรวจสอบสามเส้าด้านข้อมูล โดยพิจารณาแหล่งเวลา แหล่งสถานที่ และแหล่งบุคคลที่แตกต่างกันกล่าวคือ ถ้าข้อมูลต่างเวลากันจะเหมือนกันหรือไม่ ถ้าข้อมูลต่างสถานที่จะเหมือนกันหรือไม่ และถ้าบุคคลผู้ให้ข้อมูลเปลี่ยนไปข้อมูลจะเหมือนเดิมหรือไม่ 2) การตรวจสอบสามเส้าด้านผู้วิจัย โดยการเปลี่ยนตัวผู้สังเกตหรือสัมภาษณ์ และ 3) การตรวจสอบสามเส้าด้านวิธีรวบรวมข้อมูล โดยใช้วิธีเก็บรวบรวมข้อมูลต่าง ๆ กันเพื่อรวบรวมข้อมูลเรื่องเดียวกัน เช่น ใช้วิธีสังเกตควบคู่ไปกับการซักถาม

6. การวิจัยเชิงปริมาณ วิเคราะห์ข้อมูลโดยอาศัยคณิตศาสตร์หรือสถิติขั้นสูงด้วยการป้อนข้อมูลลงในเครื่องคอมพิวเตอร์ ซึ่งในปัจจุบันมักจะนิยมใช้โปรแกรมสำเร็จรูป เช่น SPSS แต่การวิจัยเชิงคุณภาพเป็นการวิเคราะห์โดยการตีความในคำพูด ความรู้สึกหรือความคิดเห็นของคนที่เกี่ยวข้อง โดยโยงไปถึงแนวคิดทฤษฎีเพื่อให้ความหมายแก่ข้อมูลที่ได้ หลังจากนั้นจึงทำการสร้างข้อสรุปในเรื่องดังกล่าวขึ้น แต่หากเป็นการศึกษาเอกสาร (Documentary Research) การวิเคราะห์ข้อมูลจะใช้วิธีที่เรียกว่า การวิเคราะห์เนื้อหา (Content Analysis) ให้เห็นว่าใคร ทำอะไร ที่ไหน เมื่อไร อย่างไร...

7. การวิจัยเชิงปริมาณ นำเสนอข้อมูลด้วยตัวเลขในลักษณะของตารางประการบรรยาย มีความกระชับตรงไปตรงมาตามหลักฐานหรือผลการวิเคราะห์ข้อมูล ส่วนการวิจัยเชิงคุณภาพนำเสนอข้อมูลโดยการพรรณนาความเชิงบอกเล่าในลักษณะของตัวอักษรและรูปภาพของคนในสถานการณ์ที่ศึกษา เสริมด้วยข้อมูลซึ่งเป็นคำพูดที่น่าสนใจของให้ข้อมูล นอกจากนี้อาจกล่าวถึงบริบทในสภาพแวดล้อม โดยมีข้อมูลตัวเลขประกอบบ้าง เพื่อความน่าเชื่อถือหนักแน่นของการวิจัยก็ได้

สรุปได้ว่า การวิจัยทั้ง 2 วิธีนี้มีความแตกต่างกัน การวิจัยเชิงคุณภาพมีความยืดหยุ่นกว่าการวิจัยเชิงปริมาณ ทั้งในเรื่องความเป็นธรรมชาติในวิธีการศึกษา สามารถเข้าถึงข้อมูลได้อย่างผสมกลมกลืนกับบริบท ความเป็นองค์รวมของการศึกษาโดยไม่แยกส่วน วิธีการเก็บรวบรวมข้อมูล และจำนวนผู้ให้ข้อมูล แต่ก็มีข้อจำกัดในเรื่องของเวลาและทำได้เฉพาะกรณี ซึ่งไม่สามารถทำการศึกษากับคนจำนวนมากได้หรือไม่ใช่ตัวแทนของประชากรที่ศึกษาทั้งหมด ดังนั้นข้อค้นพบจึงไม่สามารถนำไปใช้อ้างอิงกับกรณีอื่น ๆ ได้หรือได้น้อย อีกทั้งการศึกษาในภาคสนามตัวผู้วิจัยถือเป็นเครื่องมือสำคัญในการเก็บรวบรวมข้อมูล โดยจะต้องได้รับการฝึกฝนมาเป็นพิเศษสำหรับการวิจัยชนิดนี้ จึงมักถูกโจมตีในเรื่องอคติส่วนตัวหรือความรู้สึกนึกคิดที่หล่อหลอมมาจากประสบการณ์เดิม ซึ่งอาจส่งผลโน้มเอียงต่อความเที่ยงตรงและความน่าเชื่อถือของข้อมูล ดังนั้นการจะเลือกใช้วิธีการวิจัยเชิงปริมาณหรือการวิจัยเชิงคุณภาพนั้น ผู้วิจัยจะต้องตระหนักถึงเนื้อหาความรู้ในศาสตร์หรือสาขาวิชานั้น ๆ รวมถึงลักษณะของเรื่องที่จะทำการวิจัย โจทย์หรือคำถามการวิจัย วัตถุประสงค์ในการวิจัย ตลอดจนความถนัดหรือความสนใจของผู้วิจัยด้วย

3. บทบาทของสถิติในการวิจัยทางสังคมศาสตร์

การวิจัยทางสังคมศาสตร์และการบริหารนิยมใช้ทั้ง สถิติเชิงพรรณนา (Descriptive Statistics) และ สถิติเชิงอนุมาน (Inferential Statistics) ในกระบวนการวิจัย ซึ่งสามารถสรุปเป็น 3 ขั้นตอนใหญ่ ๆ ดังนี้

1) ขั้นตอนการเลือกตัวอย่างและกำหนดขนาดตัวอย่าง (Sampling)

เป็นการนำ สถิติเชิงอนุมานแบบมีพารามิเตอร์ มาใช้ โดยอาศัยหลักทฤษฎีการสุ่มตัวอย่าง (Sampling Theory) เพื่อเลือกกลุ่มตัวอย่างที่เป็นตัวแทนของประชากรอย่างเหมาะสม และใช้หลัก **ความน่าจะเป็น** รวมถึงวิธีการทางสถิติในการกำหนดขนาดตัวอย่าง (Sample Size Determination) ให้มีความเพียงพอและเป็นไปตามเกณฑ์ความเชื่อมั่น (Confidence Level) ที่ต้องการ เช่น ตารางของ **Krejcie & Morgan** หรือสูตรการคำนวณจากทฤษฎีความน่าจะเป็น

2) ขั้นตอนการบรรยายลักษณะข้อมูล (Describing Data)

เป็นการนำ สถิติเชิงพรรณนา มาใช้เพื่ออธิบายลักษณะและคุณสมบัติของข้อมูลที่รวบรวมมา โดยแสดงให้เห็นถึงแนวโน้มและการกระจายของข้อมูล ตัวอย่างเช่น

- การแจกแจงความถี่ (Frequency Distribution)
- การหาค่าร้อยละ (Percentage)
- การหาค่ากลาง เช่น ค่าเฉลี่ย (Mean) ค่ามัธยฐาน (Median) ค่าฐานนิยม (Mode)

- การวัดการกระจาย เช่น ค่าเบี่ยงเบนมาตรฐาน (Standard Deviation) พิสัย (Range) ความแปรปรวน (Variance)

3) ขั้นตอนการหาข้อสรุปจากข้อมูลตัวอย่าง (Drawing Conclusions from Sample Data)

เป็นการใช้ สถิติเชิงอนุมาน เพื่ออธิบายและสรุปผลจากข้อมูลตัวอย่างไปยังประชากร โดยเทคนิคที่สำคัญ ได้แก่

(1) การประมาณค่า (Estimation)

ใช้ค่าที่ได้จากตัวอย่างเพื่อประมาณค่าของประชากร เช่น

- ประมาณค่าใช้จ่ายเฉลี่ยของนักศึกษาในหนึ่งภาคเรียน
- ประมาณรายได้เฉลี่ยต่อวันของห้างสรรพสินค้า
- ประมาณปริมาณการใช้น้ำมันเฉลี่ยต่อวันของประชาชนไทย

การประมาณค่ามีทั้งแบบจุดประมาณ (Point Estimate) และช่วงประมาณ (Interval Estimate)

(2) การทดสอบสมมติฐาน (Hypothesis Testing)

นำข้อมูลจากตัวอย่างมาตรวจสอบข้อสมมติที่ตั้งไว้ เช่น

- การทดสอบสมมติฐานว่า อายุขัยเฉลี่ยของคนไทยไม่เกิน 80 ปี
- การทดสอบว่าสัดส่วนผู้ป่วยโรคเอดส์ในแต่ละภูมิภาคไม่แตกต่างกันอย่างมีนัยสำคัญ

(3) การหาความสัมพันธ์ (Correlation & Association)

เป็นการศึกษาความสัมพันธ์ระหว่างตัวแปรตั้งแต่ 2 ตัวขึ้นไป เช่น

- ความสัมพันธ์ระหว่างค่าใช้จ่ายกับอายุและชั้นปีของนักศึกษา
- ความสัมพันธ์ระหว่างอาชีพกับเพศของประชาชน

(4) การพยากรณ์ (Forecasting & Prediction)

เป็นการวิเคราะห์เพื่อสร้างแบบจำลองทางสถิติ (Statistical Model) เพื่อนำไปใช้ทำนายหรือพยากรณ์ข้อมูลในอนาคต เช่น

- การใช้สมการถดถอยเชิงเส้น (Linear Regression) เพื่อพยากรณ์รายได้
- การสร้างแบบจำลองอนุกรมเวลา (Time Series Model) เพื่อทำนายยอดขาย

ระดับของเทคนิคทางสถิติที่ใช้ในการวิจัย

1. เทคนิคทางสถิติขั้นพื้นฐาน (Basic Statistics)

เช่น การวัดแนวโน้มเข้าสู่ส่วนกลาง การวัดการกระจาย การหาความสัมพันธ์เชิงสหสัมพันธ์ (Correlation) และการวิเคราะห์การถดถอย (Regression Analysis)

2. เทคนิคทางสถิติขั้นสูง (Advanced Statistics)

เช่น การวิเคราะห์องค์ประกอบ (Factor Analysis) การวิเคราะห์จำแนก (Discriminant Analysis) การวิเคราะห์กลุ่ม (Cluster Analysis) การวิเคราะห์เส้นทาง (Path Analysis) หรือการสร้างแบบจำลองสมการโครงสร้าง (Structural Equation Modeling: SEM)

สรุปได้ว่า บทบาทของสถิติในการวิจัยทางสังคมศาสตร์มีความสำคัญอย่างยิ่ง เพราะทำหน้าที่เป็น เครื่องมือในการเก็บ วิเคราะห์ และแปลผลข้อมูล เพื่อให้ได้มาซึ่งข้อสรุปที่มีความถูกต้อง น่าเชื่อถือ และสามารถนำไปใช้ประกอบการตัดสินใจเชิงวิชาการ และเชิงนโยบายได้อย่างเหมาะสม โดยสถิติสามารถช่วยนักวิจัยได้ตั้งแต่การออกแบบการวิจัย การกำหนดตัวอย่าง การบรรยายข้อมูล ไปจนถึงการวิเคราะห์เพื่อสรุปผลเชิงอนุมาน ทั้งนี้ สถิติยังแบ่งได้เป็น สถิติขั้นพื้นฐาน ที่ใช้ในงานทั่วไปเพื่ออธิบายและตรวจสอบความสัมพันธ์ และ สถิติขั้นสูง ที่ใช้สร้างแบบจำลองและวิเคราะห์เชิงลึกเพื่ออธิบายปรากฏการณ์ที่ซับซ้อน ดังนั้น การเลือกใช้สถิติที่เหมาะสมกับปัญหาวิจัย วัตถุประสงค์ และลักษณะของข้อมูล จึงเป็นสิ่งจำเป็นที่ช่วยให้งานวิจัยทางสังคมศาสตร์มีคุณภาพ มีความถูกต้องตามหลักวิชาการ และสามารถสร้างองค์ความรู้ใหม่ที่มีคุณค่าในการพัฒนาสังคมได้อย่างแท้จริง

4. กระบวนการวิจัยและตำแหน่งของสถิติ

การวิจัยมักเริ่มต้นจากการวางแผนงานวิจัยและตามด้วยขั้นตอนต่าง ๆ ตั้งแต่เริ่มต้นจนกระทั่งเสร็จสิ้นการวิจัยเพื่อให้การวิจัยเป็นไปในทิศทางที่ผู้วิจัยกำหนด โดยในทุกขั้นตอนจะมีสถิติเข้ามาเกี่ยวข้อง ดังนี้

1) การวางแผน (Planning) ก่อนเริ่มการทำวิจัยไม่ว่าจะขั้นตอนใดก็ตาม ผู้วิจัยจะต้องวางแผนและกำหนดทิศทางของการวิจัย เพื่อการวิจัยและผลการศึกษาที่มีคุณภาพสามารถตอบคำถามและวัตถุประสงค์การวิจัยได้อย่างสมบูรณ์ และเพื่อก่อให้เกิด

งานวิจัยที่น่าเชื่อถือ ส่งผลต่อการรับรู้ถึงคุณภาพหรือประโยชน์และการนำไปใช้งานหรือพัฒนาองค์กร หรือแม้แต่การต่อยอดในเชิงวิชาการซึ่งในขั้นตอนนี้ สถิติจะถูกดึงเข้ามาเพื่อช่วยเกี่ยวกับข้อมูลต่างๆ ๆ ที่เป็นส่วนสนับสนุนการวิจัยตามที่ต้องการ รวมถึงการวางแผนและกำหนดประเภทของสถิติที่เหมาะสมในการวิจัย ดังนั้น สถิติจึงมีบทบาทและส่วนเกี่ยวข้องกับการวิจัยตั้งแต่เริ่มต้นการวิจัยจนกระทั่งสิ้นสุดการวิจัย

2) การออกแบบ (Design) เป็นส่วนที่สำคัญส่วนหนึ่งในการวิจัย เพราะการออกแบบการวิจัยที่ดีต้องสอดคล้องกับทุกส่วนของการวิจัย ทั้งวัตถุประสงค์ คำถามการวิจัยเกณฑ์การเลือกกลุ่มตัวอย่าง เครื่องมือการวิจัย การทดสอบเครื่องมือ การวิเคราะห์ผลสถิติ การเขียนรายงานผล และการอภิปรายผล ทั้งหมดต้องสอดคล้องกันภายใต้ข้อกำหนดของวิธีวิจัย ดังนั้นการออกแบบงานวิจัยนั้น จึงมีสถิติเป็นเครื่องมือที่เป็นตัวสนับสนุนและช่วยให้การออกแบบงานวิจัย ให้สอดคล้องกับงานวิจัย

3) การเก็บรวบรวมข้อมูล (Data Collection) ในขั้นตอนนี้สถิติมีส่วนร่วมในเรื่องการกำหนดวิธีการเก็บรวบรวม ข้อมูลที่จะนำมาซึ่งข้อมูลที่เป็นตัวแทนที่ดีของประชากร และจำนวนข้อมูลถูกต้องครบถ้วน และเป็นตัวแทนที่สามารถตอบคำถามและวัตถุประสงค์งานวิจัยนั้นได้

4) กระบวนการจัดการข้อมูล (Data Processing) คือการจัดเก็บข้อมูลที่ได้จากการเก็บรวบรวมให้อยู่ในรูปแบบที่สะดวก เข้าใจง่าย ไม่ซับซ้อน และนำไปใช้ได้อย่างรวดเร็ว โดยเริ่มตั้งแต่กระบวนการตรวจสอบความถูกต้องครบถ้วนของข้อมูลที่ได้จากเครื่องมือการวิจัย จากนั้นบันทึกข้อมูลโดยใช้โปรแกรมสำเร็จรูปทางสถิติ เช่น PSPP, Excel, หรือ SPSS โดยเก็บข้อมูลตามรูปแบบและหมวดหมู่ของตัวแปร โดยการตรวจสอบความสอดคล้องและความถูกต้องของข้อมูล สามารถนำเข้าสู่ขั้นตอนการทดสอบทางสถิติ ดังนั้นหากข้อมูลไม่มีความเหมาะสม หรือขาดคุณภาพในการเลือกใช้สถิติที่ต้องการใช้วิจัย ก็จะส่งผลต่อการวิเคราะห์ข้อมูลและความถูกต้องและมาตรฐานของการวิจัย

5) การวิเคราะห์ข้อมูล (Data Analysis) ขั้นตอนนี้สถิติมีบทบาทสำคัญอย่างยิ่งเพราะถือเป็นเครื่องมือหลัก เนื่องจากการวิเคราะห์ข้อมูลนั้นจะต้องพิจารณาเลือกสถิติที่เหมาะสมในการวิเคราะห์ข้อมูลที่ได้จากการเก็บรวบรวม เพื่อสามารถตอบวัตถุประสงค์และคำถามงานวิจัยได้อย่างเหมาะสม หากข้อมูลที่ได้มีความถูกต้องและมีคุณภาพ หากแต่การเลือกใช้สถิติที่ไม่มีความเหมาะสมมาใช้วิเคราะห์ข้อมูล ย่อมส่งผลให้งานวิจัยขาดความน่าเชื่อถือของงานวิจัยและผู้วิจัยด้วย โดยการเลือกสถิติที่เหมาะสมกับการวิเคราะห์ข้อมูล ต้องคำนึงถึงวัตถุประสงค์การวิจัย ตัวแปรที่ต้องการศึกษาตามกรอบแนวคิดการวิจัย (ตัวแปรตามหรือตัวแปรผล) กลุ่มประชากรและตัวอย่าง และเครื่องมือที่ใช้ในการศึกษาวิจัย

ดังนั้นก่อนวิเคราะห์ข้อมูลควรพิจารณาความเหมาะสมตามองค์ประกอบการวิจัย อีกทั้งเลือกสถิติให้เหมาะสมกับงานวิจัยที่ศึกษา เพื่อก่อให้เกิดคุณภาพและน่าเชื่อถือของงานวิจัย

6) การนำเสนอ (Presentation) การนำเสนอข้อมูลที่ได้จากการวิเคราะห์ข้อมูล อาจจะอยู่ในรูปของ การบรรยาย ตาราง กราฟ มีการจัดหมวดหมู่ของตัวเลขข้อมูลต่าง ๆ ที่ได้จากการเก็บรวบรวมข้อมูล เพื่อให้ผู้อ่านเข้าใจได้ง่าย นั่นคือการนำสถิติเข้ามาใช้ในขั้นตอนนี้ แต่ถ้าการวิเคราะห์ข้อมูลถูกต้องแต่มีการนำเสนอที่ไม่ถูกต้องอาจจะทำให้ผู้อ่านขาดข้อมูลที่ถูกต้องบางอย่างไปได้

7) การแปลผล (Interpretation) เมื่อถึงขั้นตอนนี้ซึ่งเป็นขั้นตอนที่สำคัญที่สุดของการทำวิจัย หากแต่เกิดความผิดพลาดในการแปลผลการวิจัยผิดนั้น จะไม่เกิดประโยชน์ใด ๆ ดังนั้นสถิติจึงเข้ามาเพื่อใช้เป็นเครื่องมือที่ช่วยเหลือในเรื่องของการแปลผลการวิจัยเพื่อให้นักวิจัยนั้นมีคุณค่า น่าเชื่อถือ สร้างความมั่นใจต่อการนำไปใช้งานหรือพัฒนาได้อย่างมีประสิทธิภาพ

8) การเผยแพร่ตีพิมพ์ (Publication) การตีพิมพ์เผยแพร่สิ่งที่ได้ค้นพบใหม่ไปประยุกต์ใช้ หรือเพิ่มพูนความรู้ให้มากขึ้น ดังนั้นในการเผยแพร่ตีพิมพ์เอกสารสถิติจึงเข้ามามีส่วนเกี่ยวข้องกับการวิจัย โดยนักสถิติจะเข้ามาให้คำแนะนำเกี่ยวกับการนำเสนอสถิติในการตีพิมพ์เผยแพร่ และการแนะนำการนำเสนอข้อมูลสถิติที่เหมาะสม เป็นต้น (ณัฐรฐนนท์ กานต์วีกุลธนา, 2565)

สรุปได้ว่า กระบวนการวิจัยตั้งแต่เริ่มต้นจนถึงการเผยแพร่ผลงาน ล้วนมี สถิติเข้ามามีบทบาทอย่างต่อเนื่องและสำคัญในทุกขั้นตอน ตั้งแต่ การวางแผน (Planning) และการออกแบบวิจัย (Design) ซึ่งต้องอาศัยหลักสถิติในการกำหนดทิศทาง เลือกตัวอย่าง และออกแบบเครื่องมือที่เหมาะสม ต่อด้วยขั้นตอน การเก็บรวบรวมข้อมูล (Data Collection) และ การจัดการข้อมูล (Data Processing) ที่ต้องใช้สถิติช่วยให้ข้อมูลมีความครบถ้วน เป็นระบบ และพร้อมเข้าสู่การวิเคราะห์ เมื่อเข้าสู่ขั้นตอน การวิเคราะห์ข้อมูล (Data Analysis) สถิติจึงมีบทบาทสูงสุด เพราะเป็นเครื่องมือหลักในการตรวจสอบสมมติฐาน หาความสัมพันธ์ สรุปผล และสร้างแบบจำลองเพื่อการพยากรณ์ หากเลือกใช้สถิติที่ถูกต้องสอดคล้องกับวัตถุประสงค์ ตัวแปร และประชากร ก็จะทำให้งานวิจัยมีคุณภาพและน่าเชื่อถือ นอกจากนี้ สถิตียังเกี่ยวข้องกับการนำเสนอผลการวิจัย (Presentation) ผ่านตาราง กราฟ หรือข้อความที่เข้าใจง่าย และในขั้นตอน การแปลผล (Interpretation) สถิติก็ช่วยให้ข้อสรุปมีความแม่นยำ ตอบคำถามวิจัยได้อย่างชัดเจน รวมถึงการ เผยแพร่ตีพิมพ์ (Publication) ที่ต้องใช้การรายงานผลเชิงสถิติให้เป็นมาตรฐาน และเป็นที่ยอมรับในวงวิชาการ กล่าวโดยสรุป สถิติทำหน้าที่เป็น แกนกลางของกระบวนการวิจัย ทั้งในด้านการออกแบบ การดำเนินการ การวิเคราะห์ และการนำเสนอ

จึงเป็นปัจจัยสำคัญที่ช่วยให้งานวิจัยทางสังคมศาสตร์และการบริหารมีความถูกต้องเที่ยงตรง น่าเชื่อถือ และสามารถนำไปใช้ประโยชน์ได้อย่างแท้จริง

5. ธรรมชาติของข้อมูลทางสังคม

ธรรมชาติของข้อมูลทางสังคม คือ ข้อมูลที่สะท้อนพฤติกรรม ความสัมพันธ์ และปรากฏการณ์ทางสังคมที่เกี่ยวข้องกับมนุษย์ ซึ่งมีความซับซ้อนและเปลี่ยนแปลงอยู่ตลอดเวลา ข้อมูลเหล่านี้มีลักษณะเฉพาะตัว เช่น ความเป็นอัตวิสัย (Subjective) ที่ขึ้นอยู่กับมุมมองและประสบการณ์ของแต่ละบุคคล อีกทั้งมักเกี่ยวข้องกับข้อมูลเชิงคุณภาพ (Qualitative Data) ที่ต้องอาศัยการตีความและการวิเคราะห์เชิงลึก (สมชาย วรจิเกษมสกุล, 2554)

การจำแนกประเภทของข้อมูลทางสังคม

1) จำแนกตามศาสตร์

(1) ข้อมูลทางประวัติศาสตร์และมานุษยวิทยา

ข้อมูลหรือหลักฐานที่เกี่ยวข้องกับผลงานและผลผลิตของมนุษย์ที่สามารถตรวจสอบได้ เช่น เอกสาร สิ่งพิมพ์ หรือโบราณวัตถุ

(2) ข้อมูลทางวิทยาศาสตร์

- **ข้อมูลทางวิทยาศาสตร์กายภาพ** ข้อมูลที่ได้จากลักษณะทางกายภาพของวัตถุและปรากฏการณ์ทางธรรมชาติ เช่น ก้อนหิน แร่ธาตุ ปริมาณน้ำฝน หรือการระเบิดของภูเขาไฟ

- **ข้อมูลทางวิทยาศาสตร์ชีวภาพ** ข้อมูลที่เกี่ยวข้องกับโครงสร้างและการทำงานของอวัยวะในสิ่งมีชีวิต เช่น วัฏจักรชีวิตของแมลง การสูบบุหรี่หลอดในร่างกายนมนุษย์

(3) ข้อมูลทางสังคมศาสตร์และพฤติกรรมศาสตร์

- **พฤติกรรมภายใน** ข้อมูลเกี่ยวกับสภาวะทางจิตใจหรือการทำงานภายในร่างกายที่ไม่สามารถสังเกตได้โดยตรง เช่น สติปัญญา ความรู้สึก การรับรู้ การตัดสินใจ

- **พฤติกรรมภายนอก** ข้อมูลเกี่ยวกับพฤติกรรมที่มนุษย์แสดงออกและสามารถสังเกตหรือวัดได้

- พฤติกรรมทางวาจา การพูด การบอกเล่า

- พฤติกรรมทางกาย การเคลื่อนไหว การเปลี่ยนแปลงทางร่างกาย

- **พฤติกรรมโมลาร์** พฤติกรรมที่สังเกตได้ง่าย เช่น การร้องไห้ การเดิน การวิ่ง

- พฤติกรรมโมเลกุลาร์ พฤติกรรมที่ต้องใช้เครื่องมือช่วย
ตรวจวัด เช่น การเต้นของหัวใจ ความดันโลหิต

2) จำแนกตามแหล่งที่มาของข้อมูล

(1) ข้อมูลปฐมภูมิ (Primary Data)

ข้อมูลที่ได้จากแหล่งกำเนิดโดยตรง ซึ่งมีความเที่ยงตรงและเชื่อถือได้สูง แต่
ใช้เวลา งบประมาณ และแรงงานมาก

ข้อดี

1. สอดคล้องกับวัตถุประสงค์การวิจัย
2. สามารถควบคุมคุณภาพการเก็บรวบรวมได้
3. ได้ข้อมูลเชิงลึกและตรงประเด็น

ข้อจำกัด

1. ใช้เวลานาน
2. สิ้นเปลืองงบประมาณ
3. หากผู้วิจัยขาดทักษะ อาจได้ข้อมูลไม่ครบถ้วนและขาดความน่าเชื่อถือ

(2) ข้อมูลทุติยภูมิ (Secondary Data)

ข้อมูลที่ได้จากการเก็บรวบรวมไว้แล้ว เช่น รายงาน หนังสือ หรือสถิติของ
หน่วยงาน ซึ่งใช้เวลาและค่าใช้จ่ายน้อยกว่าข้อมูลปฐมภูมิ

ข้อดี

1. ประหยัดเวลา งบประมาณ และแรงงาน
2. สามารถใช้ศึกษาการเปลี่ยนแปลงหรือแนวโน้มทางประวัติศาสตร์

ข้อจำกัด

1. อาจไม่สอดคล้องกับวัตถุประสงค์การวิจัย
2. ขาดความสมบูรณ์และความน่าเชื่อถือ
3. อาจไม่ทันสมัย สอดคล้องกับบริบทปัจจุบัน
4. ต้องใช้ความรอบคอบในการตีความและวิเคราะห์

3) จำแนกตามสภาพที่เป็นจริง

(1) ข้อเท็จจริง (Fact) ข้อมูลที่ไม่เปลี่ยนแปลง เช่น วันเดือนปีเกิด น้ำหนัก
ความสูง

(2) เจตคติ (Attitude) ข้อมูลที่สะท้อนทัศนคติหรือความพึงพอใจ อาจ
เปลี่ยนแปลงตามสถานการณ์

(3) ความคิดเห็นหรือความรู้สึก (Opinion/Feeling) ข้อมูลที่สะท้อน
การรับรู้หรือปฏิกิริยา มีความคงที่น้อยที่สุด

4) จำแนกตามคุณสมบัติของข้อมูล

(1) ข้อมูลปรนัย (Objective Data) ข้อมูลที่เก็บได้โดยตรง มีความเที่ยงตรงสูง เช่น การวัดความสูง น้ำหนัก

(2) ข้อมูลอัตนัย (Subjective Data) ข้อมูลที่ขึ้นอยู่กับความรู้สึก เช่น ระดับความพึงพอใจ

5) จำแนกตามลักษณะของข้อมูล

(1) ข้อมูลเชิงปริมาณ (Quantitative Data) ข้อมูลที่สามารถวัดเป็นตัวเลข เช่น คะแนนสอบ รายได้

(2) ข้อมูลเชิงคุณภาพ (Qualitative Data) ข้อมูลที่อธิบายลักษณะ เช่น ข้อความ ความคิดเห็น

6) จำแนกตามความเกี่ยวข้องกับกลุ่มตัวอย่าง

(1) ข้อมูลส่วนบุคคล (Personal Data) อายุ เพศ อาชีพ ศาสนา ระดับการศึกษา

(2) ข้อมูลสิ่งแวดล้อม (Environmental Data) ภูมิอากาศ ภูมิประเทศ สภาพแวดล้อมที่อยู่อาศัย

(3) ข้อมูลพฤติกรรม (Behavioral Data)

- พุทธิพิสัย (Cognitive Domain) ข้อมูลด้านสติปัญญา เช่น ผลสัมฤทธิ์ทางการเรียน

- จิตพิสัย (Affective Domain) ข้อมูลด้านจิตใจ เช่น ความสนใจ เจตคติ แรงจูงใจ

- ทักษะพิสัย (Psychomotor Domain) ข้อมูลด้านการปฏิบัติหรือทักษะ เช่น การเคลื่อนไหว การทำงาน (สมชาย วรภิรมย์สกุล, 2554)

สรุปได้ว่า ข้อมูลทางสังคมเป็นข้อมูลที่สะท้อนพฤติกรรม ความสัมพันธ์ และปรากฏการณ์ของมนุษย์ ซึ่งมีความซับซ้อน เปลี่ยนแปลงง่าย และมีลักษณะอัตวิสัย จึงต้องอาศัยทั้งข้อมูลเชิงปริมาณและเชิงคุณภาพร่วมกัน การจำแนกทำได้หลายเกณฑ์ เช่น ตามศาสตร์ แหล่งที่มา สภาพที่เป็นจริง คุณสมบัติ ลักษณะ และความสัมพันธ์กับกลุ่มตัวอย่าง ข้อมูลเหล่านี้เป็นพื้นฐานสำคัญของการวิจัยทางสังคมศาสตร์ที่นักวิจัยต้องเลือกใช้เหมาะสมเพื่อให้ได้ผลที่ถูกต้องและน่าเชื่อถือ

6. จริยธรรมในการใช้สถิติ

จริยธรรมในการใช้สถิติ หมายถึง แนวทางการปฏิบัติที่ยึดหลักความซื่อสัตย์ ความรับผิดชอบ ความเป็นกลาง และการเคารพต่อแหล่งข้อมูลและผู้เกี่ยวข้อง เพื่อให้การใช้สถิติเป็นไปอย่างถูกต้อง เที่ยงตรง และเป็นประโยชน์ต่อสังคม โดยหลีกเลี่ยงการบิดเบือนข้อมูล การนำเสนอที่ก่อให้เกิดความเข้าใจผิด การเลือกปฏิบัติ และการละเมิดสิทธิความเป็นส่วนตัว

ในการวิจัยสาขาสังคมศาสตร์ ได้มีการกำหนด **จรรยาบรรณนักวิจัย** โดยสภาวิจัยแห่งชาติ เพื่อใช้เป็นมาตรฐานกลางที่นักวิจัยต้องปฏิบัติ (สินธะวา คามดิษฐ์, 2550) ซึ่งประกอบด้วยแนวทางสำคัญดังนี้

1) นักวิจัยต้องซื่อสัตย์และมีคุณธรรมในทางวิชาการและการจัดการ โดยต้องมีความซื่อสัตย์ต่อตนเองไม่นำผลงานของผู้อื่นมาเป็นของตน ไม่ลอกเลียนงานของผู้อื่น ต้องให้เกียรติและอ้างอิงบุคคลหรือแหล่งที่มาของข้อมูลที่ใช้ในงานวิจัย ต้องซื่อตรงต่อการแสวงหาทุนวิจัยและมีความเป็นธรรมเกี่ยวกับผลประโยชน์ที่ได้จากการวิจัย

2) นักวิจัยต้องตระหนักถึงพันธกรณีในการทำงานวิจัย โดยทำตามข้อตกลงที่ทำไว้กับหน่วยงานที่สนับสนุนการวิจัยและต่อหน่วยงานที่ต้นสังกัด นักวิจัยต้องปฏิบัติตามพันธกรณีและข้อตกลงการวิจัยที่ผู้เกี่ยวข้องทุกฝ่ายยอมรับร่วมกัน อุทิศเวลาทำงานวิจัยให้ได้ผลดีที่สุด และเป็นไปตามกำหนดเวลา มีความรับผิดชอบไม่ละทิ้งงานระหว่างดำเนินการ

3) นักวิจัยต้องมีพื้นฐานความรู้ในสาขาวิชาที่ทำวิจัย โดยต้องมีพื้นฐานความรู้ในสาขาวิชาที่ทำวิจัยอย่างเพียงพอ และมีความรู้ความชำนาญ หรือมีประสบการณ์เกี่ยวเนื่องกับเรื่องที่ทำวิจัย เพื่อนำไปสู่งานวิจัยที่มีคุณภาพและเพื่อป้องกันปัญหาการวิเคราะห์ การตีความหรือการสรุปที่ผิดพลาด อันอาจก่อให้เกิดความเสียหายต่องานวิจัย

4) นักวิจัยต้องมีความรับผิดชอบต่อสิ่งที่ศึกษาวิจัย ทั้งนี้ไม่ว่าจะเป็นสิ่งที่มีชีวิตหรือไม่มีชีวิต นักวิจัยต้องดำเนินการด้วยความรอบคอบระมัดระวัง และเที่ยงตรงในการทำวิจัยที่เกี่ยวข้องกับคน สัตว์ พืช ศิลปวัฒนธรรม ทรัพยากร และสิ่งแวดล้อม มีจิตสำนึกและมีปณิธานที่จะอนุรักษ์ศิลปวัฒนธรรม ทรัพยากร และสิ่งแวดล้อม

5) นักวิจัยต้องเคารพศักดิ์ศรีและสิทธิของมนุษย์ที่ใช้เป็นตัวอย่างในการวิจัย โดยต้องไม่คำนึงถึงผลประโยชน์ทางวิชาการจนละเลยและขาดความเคารพในศักดิ์ศรีของเพื่อนมนุษย์ ต้องถือเป็นภาระหน้าที่ที่จะอธิบายจุดมุ่งหมายของการวิจัยแก่บุคคลที่เป็นกลุ่มตัวอย่าง โดยไม่หลอกลวงหรือบีบบังคับ และไม่ละเมิดสิทธิส่วนบุคคล

6) นักวิจัยต้องมีอิสระทางความคิด ปราศจากอคติในทุกขั้นตอนของการทำวิจัย โดยต้องตระหนักว่าอคติส่วนตัวหรือความลำเอียงทางวิชาการ อาจส่งผลให้มีการบิดเบือนข้อมูลและข้อค้นพบทางวิชาการ อันเป็นเหตุให้เกิดผลเสียหายต่องานวิจัย

7) นักวิจัยพึงนำผลงานวิจัยไปใช้ประโยชน์ในทางที่ชอบ โดยเผยแพร่ผลงานวิจัยเพื่อประโยชน์ทางวิชาการและสังคม ไม่ขยายผลข้อค้นพบจนเกินความเป็นจริง และไม่ใช้ผลงานวิจัยไปทางมิชอบ

8) นักวิจัยพึงเคารพความคิดเห็นทางวิชาการของผู้อื่น โดยมีใจกว้างพร้อมที่จะเปิดเผยข้อมูลและขั้นตอนการวิจัย ยอมรับฟัง ความคิดเห็นและเหตุผลทางวิชาการของผู้อื่น และพร้อมที่จะปรับปรุงแก้ไขงานวิจัยของตนให้ถูกต้อง

9) นักวิจัยพึงมีความรับผิดชอบต่อสังคมทุกระดับ โดยมีจิตสำนึกที่จะอุทิศกำลังสติปัญญาในการทววิจัย เพื่อความก้าวหน้าทางวิชาการ เพื่อความเจริญและประโยชน์สุขของสังคมและมวลมนุษยชาติ

นอกจากหลักจรรยาบรรณของนักวิจัยแล้ว ยังมีประเด็นด้านจริยธรรมการวิจัยที่เกี่ยวข้องกับคน ซึ่งงานวิจัยทางสังคมศาสตร์จำนวนไม่น้อยที่เกี่ยวข้องกับคน ซึ่งอาจเป็นนักท่องเที่ยว ชาวบ้าน พนักงานผู้ประกอบการ หรือเจ้าหน้าที่ของรัฐ เป็นต้น โดยผู้วิจัยจำเป็นต้องเก็บรวบรวมข้อมูลจากบุคคลเหล่านั้นด้วยการสอบถามหรือการสัมภาษณ์ ดังนั้นนักวิจัยทางสังคมศาสตร์จึงควรทราบถึงประเด็นต่าง ๆ ที่เกี่ยวข้องกับการวิจัยในคนด้วยเช่นกัน (อัศวิน แสงพิกุล, 2556) หลักจริยธรรมพื้นฐานในการวิจัยทางสังคมศาสตร์ ได้แก่

1) ผู้ให้ข้อมูล (research participants) เช่น นักท่องเที่ยว ชาวบ้าน พนักงานเจ้าหน้าที่ ต้องได้รับการบอกกล่าว และให้ความยินยอมโดยสมัครใจที่จะให้ข้อมูลในเรื่องเกี่ยวกับการวิจัย

2) นักวิจัยต้องรักษาความลับ ความเป็นส่วนตัว และการปิดบังชื่อของผู้ให้ข้อมูล

3) การวิจัยต้องไม่ก่อให้เกิดความเดือดร้อนทางร่างกายและจิตใจของผู้ให้ข้อมูล

4) ผลของการวิจัยควรให้ผลดีต่อสังคม (บุปผา ศิริรัศมี และคณะ, 2544)

การปกป้องความเป็นส่วนตัวและการรักษาความลับ

เบญญา ยอดดำเนิน-แอ็ดติภจ (2553) ได้กล่าวถึงการปกป้องความเป็นส่วนตัวและการรักษาความลับของผู้ให้ข้อมูลว่าเรื่องเหล่านี้ (โดยเฉพาะงานวิจัยเชิงคุณภาพ) ถือว่าเป็นเรื่องสำคัญอย่างยิ่งที่ผู้วิจัยต้องตระหนักและระมัดระวังเป็นพิเศษ เพราะเป็นสิทธิขั้นพื้นฐานของผู้ให้ข้อมูลที่จะได้รับการปกป้อง ซึ่งการปกป้องความเป็นส่วนตัวและการรักษาความลับของผู้ให้ข้อมูลมีแนวทางดังนี้

1) การออกแบบการวิจัยที่เหมาะสม สอดคล้องกับโจทย์ และผู้เข้าร่วมวิจัยจะช่วยปกป้องความเป็นส่วนตัวและการรั่วไหลของความลับได้ เช่น การตั้งคำถาม ประเด็นที่จะซักถาม และการระมัดระวังเวลาสัมภาษณ์ เป็นต้น

2) กระบวนการขอความยินยอมเป็นกลไกสำคัญประการหนึ่งที่จะช่วยปกป้องความเป็นส่วนตัวและเป็นหลักประกันการเก็บรักษาข้อมูลเป็นความลับ การขอความยินยอมเป็นช่องทางที่จะช่วยให้ผู้ที่ให้ข้อมูลสามารถประเมินและชั่งน้ำหนักระหว่างความเสี่ยงและประโยชน์ที่พึงได้ก่อนการตัดสินใจเซ็นหนังสือแสดงเจตนายินยอมเข้าร่วมวิจัย

3) หากเป็นการเก็บข้อมูลจากกลุ่มคนที่หลากหลายและเก็บเพียงครั้งเดียว การใช้วิธีแบบนิรนามก็ย่อมทำได้ (การไม่ระบุชื่อและตัวบ่งชี้บุคคล) โดยในแบบสำรวจไม่จำเป็นต้องมีชื่อหรือข้อมูลบ่งชี้บุคคลอื่นและควรใช้สถิติวิเคราะห์ที่เป็นภาพรวม ไม่ใช่ข้อมูลรายบุคคล ซึ่งจะช่วยลดความเสี่ยงต่อการรั่วไหลของข้อมูลได้

4) หากเป็นการเก็บข้อมูลในกลุ่มคนเดียวกันตั้งแต่สองครั้งขึ้นไป การใช้วิธีแบบนิรนามอาจไม่เหมาะสม เพราะผู้วิจัยต้องเชื่อมต่อข้อมูลของกลุ่มคนในรอบแรกเข้ากับรอบสองหรือรอบที่สาม ดังนั้น ผู้วิจัยจึงต้องหาวิธีการกำหนดรหัสและข้อบ่งชี้บุคคลที่ผู้วิจัยและทีมงานเท่านั้นที่จะเข้าถึงได้ โดยผู้อื่นหากจะนำข้อมูลไปใช้ก็ใช้ได้แต่จะไม่สามารถระบุหรือเชื่อมโยงถึงผู้ให้ข้อมูลได้

5) ในกรณีเป็นการสัมภาษณ์แบบสนทนากลุ่ม ผู้วิจัยควรระบุให้ชัดเจนในข้อเสนอโครงการวิจัยและชี้แจงต่อผู้เข้าร่วมวิจัยในเรื่องการเก็บรักษาข้อมูล คนที่จะเข้าถึงข้อมูล การไม่ระบุชื่อและข้อบ่งชี้บุคคลที่เชื่อมโยงไปถึงผู้ให้ข้อมูลได้ รวมทั้งการใช้นามสมมติหรือรหัสแทน ตลอดจนควรทำข้อตกลงกับผู้เข้าร่วมวิจัยว่าข้อมูลที่อภิปรายจะไม่นำไปพูดข้างนอก

6) หากเป็นการเขียนรายงานการวิจัย ผู้วิจัยไม่ระบุสถานที่ที่ทำการศึกษาและชื่อของผู้ให้ข้อมูลหรือ ผู้ที่เกี่ยวข้องอย่างชัดเจนเฉพาะเจาะจง ควรใช้นามสมมติแทน ในกรณีที่มีการยกถ้อยคำหรือข้อความของผู้ให้ข้อมูลมาอ้างอิง ผู้วิจัยก็เพียงระบุลักษณะทั่วไปของผู้ให้ข้อมูลก็พอเพื่อให้ผู้อ่านเห็นภาพเจ้าของข้อมูล (ไม่จำเป็นต้องระบุชื่อผู้ให้ข้อมูล) เช่น ชายหรือหญิง อายุเท่าไร อาชีพอะไร เป็นต้น

7) งานวิจัยบางเรื่องการใช้แผนที่ในรายงานการวิจัยมีวัตถุประสงค์เพื่อให้ผู้อ่านเห็นภาพและรู้จักพื้นที่ที่ทำการศึกษา ทั้งนี้หากมีบางประเด็นที่อ่อนไหวหรืออาจส่งผลกระทบต่อผู้ให้ข้อมูล ผู้วิจัยไม่ควรระบุที่ตั้งหรือการเข้าถึงพื้นที่และประชากรกลุ่มที่ศึกษาได้อย่างชัดเจน เช่น การระบุชื่อถนน ชื่อชุมชนหรือหมู่บ้าน

8) ข้อมูลที่เก็บรวบรวมมาทั้งหมด ควรเก็บรักษาไว้อย่างดีโดยทีมนักวิจัยไม่ควรนำไปเผยแพร่แก่สาธารณชน

9) การวิเคราะห์ข้อมูล (หากเป็นงานวิจัยเชิงคุณภาพ) ควรกระทำในภาพรวมไม่ควรเจาะจงว่าเป็นข้อมูลจากผู้ใดคนใดคนหนึ่ง

ทั้งนี้ มีสองแนวทางหลักในการคุ้มครองข้อมูล ได้แก่

1) ข้อมูลแบบนิรนาม (Anonymity) หมายถึง การเก็บข้อมูลที่มีการปกปิดชื่อและข้อมูลของบุคคล โดยผู้วิจัยไม่สามารถทราบได้ว่าใครเป็นผู้ให้ข้อมูล ซึ่งเป็นแนวทางที่ดีที่สุดในการป้องกันการรั่วไหลของข้อมูล ตัวอย่างเช่น ผู้วิจัยอาจจำผู้อื่นในการเก็บข้อมูลเป็นต้น

2) การรักษาความลับ (Confidentiality) หมายถึง การเก็บข้อมูลที่มีการระบุชื่อหรือตัวบ่งชี้อื่นๆของบุคคล โดยผู้วิจัยจะทราบว่าข้อมูลเหล่านี้มาจากไหนหรือเป็นของใคร แต่ผู้วิจัยจะยึดกฎการรักษาความลับของข้อมูลอย่างเคร่งครัด ไม่เปิดเผยให้ผู้อื่นทราบ ตลอดจนการวิเคราะห์ผลที่แสดงในภาพรวม ไม่สะท้อนให้เห็นข้อมูลเป็นรายบุคคล รวมทั้งการทำลายแบบสอบถามหรือลบข้อมูลออกจากเทปสัมภาษณ์หลังจากงานวิจัยเสร็จสิ้น

สรุปได้ว่า จริยธรรมในการใช้สถิติและการวิจัยทางสังคมศาสตร์ คือการคุ้มครองความถูกต้อง โปร่งใส และความน่าเชื่อถือของงานวิจัย ควบคู่กับการเคารพสิทธิศักดิ์ศรี และความเป็นส่วนตัวของผู้ให้ข้อมูล นักวิจัยจึงต้องมีทั้งความรู้ ความรับผิดชอบ และจิตสำนึกทางวิชาการ เพื่อให้ผลงานมีคุณค่าและสร้างประโยชน์ต่อสังคมอย่างแท้จริง

สรุปท้ายบท

สถิติและการวิจัยเป็นกลไกพื้นฐานที่ช่วยให้นักวิชาการเข้าใจปรากฏการณ์ทางสังคมได้อย่างเป็นระบบ โดยสถิติทำหน้าที่ในการรวบรวม วิเคราะห์ และตีความข้อมูลเพื่อสร้างข้อสรุปที่มีความน่าเชื่อถือ ขณะที่การวิจัยทางสังคมศาสตร์ทำหน้าที่ค้นหาความจริงเกี่ยวกับพฤติกรรม ความคิด และความสัมพันธ์ของมนุษย์ในสังคม เพื่อสร้างองค์ความรู้ใหม่และสนับสนุนการกำหนดนโยบายสาธารณะ

สถิติมีความหมายทั้งในฐานะตัวเลขที่สะท้อนข้อเท็จจริง ศาสตร์ที่ว่าด้วยระเบียบวิธีการวิเคราะห์ข้อมูล และค่าเชิงคำนวณที่ได้จากกลุ่มตัวอย่าง โดยจำแนกได้เป็น **สถิติเชิงพรรณนา** ที่ใช้สรุปและอธิบายข้อมูล และ **สถิติเชิงอนุมาน** ที่ใช้สรุปอ้างอิงจากกลุ่มตัวอย่างไปยังประชากร ซึ่งครอบคลุมการประมาณค่า การทดสอบสมมติฐาน การวิเคราะห์ความสัมพันธ์ และการพยากรณ์ ขณะเดียวกัน การวิจัยทางสังคมศาสตร์แบ่งเป็นการวิจัยเชิงคุณภาพที่มุ่งตีความและอธิบายปรากฏการณ์ในบริบทเฉพาะ และการวิจัยเชิงปริมาณที่เน้นการพิสูจน์สมมติฐานและใช้การวิเคราะห์เชิงสถิติอย่างเป็นระบบ

บทบาทของสถิติในการวิจัยปรากฏในทุกขั้นตอน ตั้งแต่การออกแบบการวิจัย การกำหนดกลุ่มตัวอย่าง การจัดการและวิเคราะห์ข้อมูล ไปจนถึงการนำเสนอและแปลผล ซึ่งสถิติทั้งขั้นพื้นฐานและขั้นสูงช่วยเสริมความแม่นยำและความน่าเชื่อถือให้แก่งานวิจัย โดยเฉพาะอย่างยิ่งในการอธิบายความสัมพันธ์และสร้างแบบจำลองที่ซับซ้อน

นอกจากนี้ การทำวิจัยทางสังคมศาสตร์ยังต้องพิจารณาถึง **ธรรมชาติของข้อมูล** ที่มีความซับซ้อนและอหิวสัย การจำแนกข้อมูลตามศาสตร์ แหล่งที่มา คุณสมบัติ และลักษณะต่าง ๆ ช่วยให้นักวิจัยเลือกใช้วิธีการที่เหมาะสม ขณะเดียวกัน **จริยธรรมในการใช้สถิติและการวิจัย** ถือเป็นปัจจัยสำคัญที่กำหนดคุณค่าของงานวิชาการ นักวิจัยต้องยึดมั่นในความซื่อสัตย์ ความเป็นกลาง และการเคารพสิทธิและศักดิ์ศรีของผู้ให้ข้อมูล รวมทั้งปฏิบัติตามแนวทางจรรยาบรรณและมาตรฐานสากล เพื่อให้งานวิจัยที่ได้มีความถูกต้อง โปร่งใส และเป็นประโยชน์ต่อสังคมโดยรวม

กล่าวโดยสรุป สถิติและการวิจัยทางสังคมศาสตร์มิใช่เพียงเครื่องมือทางวิชาการ แต่เป็นรากฐานของการสร้างองค์ความรู้ การพัฒนานโยบาย และการขับเคลื่อนสังคมให้ก้าวหน้าอย่างยั่งยืน ทั้งนี้ ความสำเร็จของการวิจัยขึ้นอยู่กับการใช้สถิติที่ถูกต้อง การเข้าใจธรรมชาติของข้อมูล และการยึดมั่นในจริยธรรมของนักวิจัยเป็นสำคัญ

บทที่ 2

ประเภทของข้อมูลและการเก็บรวบรวมข้อมูล (TYPES OF DATA AND DATA COLLECTION METHODS)

ข้อมูล (Data) ถือเป็นหัวใจสำคัญของงานวิจัยและการวิเคราะห์ทางสถิติ เพราะเป็นวัตถุดิบพื้นฐานที่นักวิจัยใช้ในการตรวจสอบข้อเท็จจริง ตอบคำถามวิจัย และสร้างข้อสรุปที่มีนัยสำคัญ ข้อมูลที่ได้มาจะมีทั้งในเชิง ปริมาณ (Quantitative Data) และ คุณภาพ (Qualitative Data) ซึ่งแต่ละประเภทมีลักษณะ วิธีการเก็บรวบรวม ตลอดจน แนวทางการวิเคราะห์ที่แตกต่างกันออกไป การเข้าใจความหมายและประเภทของข้อมูล อย่างถูกต้องจึงมีความสำคัญยิ่ง เพราะจะช่วยให้ นักวิจัยเลือกใช้วิธีการที่เหมาะสมกับ ปัญหาวิจัย และสามารถแปลผลออกมาได้ตรงตามความเป็นจริง

นอกจากประเภทของข้อมูลแล้ว ระดับของการวัด (Scale of Measurement) ก็มีความสำคัญไม่แพ้กัน เนื่องจากระดับข้อมูลที่ต่างกันจะเป็นตัวกำหนดว่าข้อมูลนั้น สามารถนำไปวิเคราะห์ด้วยสถิติใดได้บ้าง ตั้งแต่ระดับต่ำที่สุดคือ นามบัญญัติ (Nominal Scale) ไปจนถึงระดับสูงสุดคือ อัตราส่วน (Ratio Scale) ซึ่งมีความสมบูรณ์ทาง คณิตศาสตร์และสามารถประมวลผลได้อย่างละเอียด การทำความเข้าใจเรื่องระดับของ ข้อมูลจึงเป็นพื้นฐานที่จะช่วยให้นักวิจัยเลือกใช้สถิติได้ถูกต้อง

กระบวนการเก็บรวบรวมข้อมูลก็เป็นอีกขั้นตอนหนึ่งที่สำคัญ เพราะคุณภาพ ของข้อมูลขึ้นอยู่กับวิธีการและเครื่องมือที่ใช้ หากการเก็บข้อมูลไม่มีความน่าเชื่อถือ ผลการ วิเคราะห์ก็ย่อมไม่มีความหมาย ดังนั้น การเลือกวิธีเก็บข้อมูลเชิงคุณภาพ เช่น การ สัมภาษณ์และการสังเกต หรือการเก็บข้อมูลเชิงปริมาณ เช่น การใช้แบบทดสอบและ แบบสอบถาม จึงต้องมีการวางแผน ควบคุม และตรวจสอบคุณภาพของเครื่องมือ ทั้งด้าน ความตรง (Validity) และ ความเชื่อมั่น (Reliability) อย่างรอบคอบ

สุดท้าย การกำหนดกลุ่มตัวอย่างและขนาดของตัวอย่าง ก็เป็นประเด็นสำคัญที่ ส่งผลต่อความเที่ยงตรงของการวิจัย หากกลุ่มตัวอย่างไม่เป็นตัวแทนที่ดีของประชากร ผลการวิจัยจะไม่สามารถอ้างอิงได้อย่างน่าเชื่อถือ การเลือกใช้วิธีการสุ่มที่เหมาะสมและการ คำนวณขนาดของกลุ่มตัวอย่างอย่างมีหลักเกณฑ์จึงเป็นสิ่งจำเป็น บทนี้จึงมุ่งเน้นให้ความรู้ ตั้งแต่ประเภทและระดับของข้อมูล วิธีการเก็บและตรวจสอบคุณภาพเครื่องมือ ไปจนถึง หลักการสุ่มและกำหนดขนาดตัวอย่าง เพื่อเป็นรากฐานที่มั่นคงสำหรับการทำวิจัยทาง สังคมศาสตร์และการวิเคราะห์ข้อมูลเชิงสถิติในขั้นต่อไป

1. ประเภทของข้อมูล

คำว่า “ข้อมูล” ได้รับการนิยามและอธิบายโดยหลายแหล่ง ทั้งในเชิงวิชาการและการประยุกต์ใช้ เช่น ข้อมูล (Data) หมายถึง ค่าของตัวแปรทั้งเชิงคุณภาพและเชิงปริมาณที่อยู่ในความควบคุมของสิ่งต่าง ๆ ในด้านคอมพิวเตอร์ ข้อมูลปรากฏในรูปแบบโครงสร้าง เช่น ตาราง (แถวและหลัก) ต้นไม้ (ความสัมพันธ์เชิงลำดับชั้น) หรือกราฟ (ความสัมพันธ์แบบเครือข่าย) โดยทั่วไป ข้อมูลเป็นผลจากการวัดและสามารถนำเสนอผ่านกราฟหรือภาพ ข้อมูลยังหมายถึง ข้อเท็จจริงเกี่ยวกับตัวแปรที่สำรวจด้วยวิธีการวัด (สุชาติ ประสิทธิ์รัฐสินธุ์, 2544) หรือข้อเท็จจริงที่เก็บจากตัวอย่างหรือประชากร (สุจรรยา ทรัพย์ศิริโสภา, 2555) อีกทั้งยังอาจหมายถึง ข้อเท็จจริงที่ได้จากการเก็บรวบรวมคุณลักษณะจากกลุ่มเป้าหมายตามวัตถุประสงค์ที่กำหนด ทั้งข้อมูลเชิงปริมาณ เช่น คะแนน และข้อมูลเชิงคุณภาพ เช่น พฤติกรรม (สมชาย วรภิเษมสกุล, 2554) ขณะเดียวกัน ข้อมูลอาจหมายถึง ข้อเท็จจริงหรือเหตุการณ์ในรูปของตัวเลข ข้อความ สัญลักษณ์ ที่เกิดจากการวัด การสังเกต หรือการบันทึก ซึ่งต้องมีความถูกต้องและแม่นยำ (วรรณลักษณ์ เมียนเกิด, 2555)

โดยสรุป “ข้อมูล” คือ ข้อเท็จจริงหรือเรื่องราวที่เกี่ยวข้องกับสิ่งต่าง ๆ เช่น คน สัตว์ สิ่งของ หรือสถานที่ อยู่ในรูปแบบที่เหมาะสมต่อการสื่อสาร การตีความ และการประมวลผล ข้อมูลอาจเกิดจากการสังเกต การรวบรวม หรือการวัด และอยู่ได้ทั้งในรูปตัวเลข ข้อความ ภาพ เสียง หรือสัญลักษณ์อื่น ๆ โดยมีคุณสมบัติสำคัญคือ ต้องเป็นจริง แม่นยำ และต่อเนื่อง

1) ข้อมูลแบ่งตามลักษณะของข้อมูล

สุจรรยา ทรัพย์ศิริโสภา (2555) ได้แบ่งข้อมูลตามลักษณะออกเป็น 2 ประเภทใหญ่ ได้แก่ ข้อมูลเชิงปริมาณ และ ข้อมูลเชิงคุณภาพ ดังนี้

(1) ข้อมูลเชิงปริมาณ (Quantitative Data)

ข้อมูลเชิงปริมาณ คือ ข้อมูลที่สามารถวัดออกมาเป็นตัวเลขได้ และสามารถเปรียบเทียบความมาก-น้อยในเชิงปริมาณได้อย่างชัดเจน เช่น รายได้ อายุ ส่วนสูง หรือจำนวนสินค้า เป็นต้น ทั้งนี้ แบ่งออกเป็น 2 ลักษณะย่อย คือ

- ข้อมูลแบบต่อเนื่อง (Continuous Data) ข้อมูลที่มีค่าต่อเนื่องในช่วงที่กำหนด สามารถมีค่าระหว่างตัวเลขได้ เช่น ความสูง อายุ ระยะทาง

- ข้อมูลแบบไม่ต่อเนื่อง (Discrete Data) ข้อมูลที่อยู่ในรูปของจำนวนนับหรือจำนวนเต็ม ไม่สามารถแบ่งย่อยได้ เช่น จำนวนนักศึกษา จำนวนสมาชิกในครัวเรือน

วรรณลักษณ์ เมียนเกิด (2555) อธิบายว่าข้อมูลเชิงปริมาณคือข้อมูลที่อยู่ในรูปตัวเลขที่สามารถนำไปใช้วิเคราะห์ทางสถิติได้ ขณะที่อีกแนวคิดมองว่าเป็นข้อมูลที่ได้จาก

การสังเกตหรือการวัดในเชิงตัวเลข เช่น คะแนนสอบ ผลสัมฤทธิ์ทางการเรียน ความถนัดทางการเรียน หรือคุณลักษณะด้านจิตพิสัยที่สามารถวัดผลออกมาเป็นคะแนนได้

(2) ข้อมูลเชิงคุณภาพ (Qualitative Data)

ข้อมูลเชิงคุณภาพ คือ ข้อมูลที่ไม่สามารถระบุได้ในเชิงปริมาณว่ามากหรือน้อย แต่แสดงลักษณะหรือคุณสมบัติของสิ่งที่ศึกษา อาจแทนด้วยตัวเลขได้ แต่ตัวเลขนั้นไม่มีความหมายเชิงปริมาณ เช่น เพศ ระดับการศึกษา อาชีพ หรือทัศนคติ

ผู้เชี่ยวชาญได้ให้ความหมายเพิ่มเติม เช่น วรรณลักษณ์ เมียนเกิด (2555) ระบุว่าข้อมูลที่อยู่ในรูปของถ้อยคำ ข้อความ หรือคำให้สัมภาษณ์ ซึ่งผู้วิจัยได้จากการสอบถามหรือการบันทึกจากการสังเกต ส่วนทวีศักดิ์ นพเกษร (2548) อธิบายว่าข้อมูลเชิงคุณภาพครอบคลุมตั้งแต่ความคิด ความเห็น ความเชื่อ เจตคติ คุณค่า โลกทัศน์ พฤติกรรม วิถีชีวิต ปฏิสัมพันธ์ กระบวนการทางสังคมภายในกลุ่มหรือองค์กร ตลอดจนการรับรู้ (Perception) อารมณ์ และความรู้สึกของมนุษย์

2) ข้อมูลแบ่งตามแหล่งที่มา

ข้อมูลสามารถจำแนกได้ 2 ประเภทหลักตามแหล่งที่มา ได้แก่ **ข้อมูลปฐมภูมิ** และ **ข้อมูลทุติยภูมิ** ดังนี้

(1) ข้อมูลปฐมภูมิ (Primary Data)

ข้อมูลที่ผู้วิจัยหรือผู้ใช้เป็นผู้เก็บรวบรวมขึ้นเองโดยตรงจากแหล่งข้อมูลหรือกลุ่มเป้าหมาย เช่น การสำรวจ การสัมภาษณ์ การสังเกต หรือการทดลอง ข้อมูลปฐมภูมิจึงมีความทันสมัย ถูกต้อง และน่าเชื่อถือสูง เนื่องจากเก็บขึ้นมาใหม่และสอดคล้องกับวัตถุประสงค์ที่ต้องการ อย่างไรก็ตาม ข้อมูลประเภทนี้มีข้อจำกัดที่สำคัญ ได้แก่ ใช้เวลาและงบประมาณมาก ต้องใช้กำลังคนจำนวนมาก และหากผู้วิจัยขาดความชำนาญในการเก็บข้อมูล อาจทำให้ข้อมูลไม่ครบถ้วนหรือคลาดเคลื่อนได้

(2) ข้อมูลทุติยภูมิ (Secondary Data)

ข้อมูลที่ได้จากการเก็บรวบรวมไว้แล้ว โดยผู้นำไปใช้ไม่ได้เป็นผู้เก็บเอง แต่ดึงมาจากแหล่งข้อมูลเดิม เช่น รายงานการวิจัย หนังสือ วารสาร บทความ สถิติทางการ หรือฐานข้อมูลออนไลน์ ข้อมูลประเภทนี้มีข้อดีคือ ประหยัดเวลา แรงงาน และค่าใช้จ่าย อีกทั้งสามารถใช้เพื่อศึกษาการเปลี่ยนแปลงหรือแนวโน้มทางประวัติศาสตร์ได้ แต่ก็มีข้อจำกัด เช่น อาจไม่สอดคล้องกับวัตถุประสงค์เฉพาะของงานวิจัย ขาดรายละเอียดหรือความถูกต้องครบถ้วน และอาจล้าสมัยไม่ทันต่อสถานการณ์ปัจจุบัน

สรุปได้ว่า ข้อมูล หมายถึง ข้อเท็จจริงหรือเรื่องราวเกี่ยวกับสิ่งต่าง ๆ ที่อยู่ในรูปแบบซึ่งสามารถใช้เพื่อการสื่อสาร การวิเคราะห์ และการประมวลผลได้ ข้อมูลแบ่งได้หลายประเภท ทั้งตามลักษณะ (เชิงปริมาณและเชิงคุณภาพ) และตามแหล่งที่มา (ปฐมภูมิ

และทฤษฎี) ซึ่งแต่ละประเภทย่อมมีข้อดีและข้อจำกัด การเลือกใช้ข้อมูลจึงต้องสอดคล้องกับวัตถุประสงค์และลักษณะของงานวิจัยเพื่อให้ได้ผลลัพธ์ที่ถูกต้อง น่าเชื่อถือ และนำไปใช้ประโยชน์ได้จริง

2. ระดับของข้อมูล

สมชาย วรภิรมย์สกุล (2554) อธิบายว่า ข้อมูลในการวิจัยจำนวนมากได้มาจากการวัด (Measurement) ซึ่งหมายถึง การกำหนดตัวเลขหรือสัญลักษณ์แทนปริมาณ คุณภาพ หรือคุณลักษณะของสิ่งที่ต้องการศึกษา เพื่อนำข้อมูลเชิงปริมาณที่ได้ไปวิเคราะห์และสรุปผล โดยข้อมูลสามารถจำแนกออกเป็น 4 ระดับ ได้แก่ นามบัญญัติ เรียงลำดับ ช่วง และอัตราส่วน

1) ข้อมูลระดับนามบัญญัติ (Nominal Scale)

ข้อมูลระดับนามบัญญัติ คือ ข้อมูลที่ใช้จำแนกหรือจัดกลุ่มสิ่งต่าง ๆ ตามชื่อหรือคุณลักษณะ โดยไม่มีความหมายเชิงปริมาณ และไม่สามารถนำไปคำนวณทางคณิตศาสตร์ เช่น บวก ลบ คูณหาร หรือเปรียบเทียบเชิงปริมาณได้ ตัวเลข (ถ้ามี) เป็นเพียง “รหัส” ที่แทนกลุ่มเท่านั้น เช่น กำหนดให้เพศชาย = 1 และเพศหญิง = 2 ซึ่งตัวเลขนี้ไม่ใช่ค่าที่แสดงความมาก-น้อย

ตัวอย่างข้อมูลระดับนามบัญญัติ

- เพศ ☐ ชาย ☐ หญิง
- อาชีพ ☐ รับราชการ ☐ รัฐวิสาหกิจ ☐ รับจ้าง ☐ อาชีพส่วนตัว
- ศาสนา ☐ พุทธ ☐ คริสต์ ☐ อิสลาม
- หมายเลขโทรศัพท์ 000-0000000

ลักษณะสำคัญ (บุญธรรม กิจปรีดาบริสุทธิ์, 2534)

1. ข้อมูลอยู่แยกกลุ่มกันอย่างไม่เป็นลำดับ
2. แต่ละกลุ่มไม่มีความสัมพันธ์กัน
3. การใช้ตัวเลขหรือสัญลักษณ์แทนกลุ่มไม่ได้มีความหมายเชิงปริมาณ
4. สมาชิกในกลุ่มเดียวกันมีคุณลักษณะเหมือนกัน

ค่าสถิติที่ใช้ได้ ความถี่ (Frequency) ร้อยละ (Percentage) ฐานนิยม (Mode)

สรุปได้ว่า หากข้อมูลบอกได้เพียงความแตกต่างในลักษณะชื่อ เช่น ชื่อนักเรียนที่แตกต่างกัน แต่ไม่สามารถบอกปริมาณหรือคุณภาพอื่น ๆ ได้ ข้อมูลนั้นจัดเป็นข้อมูลระดับนามบัญญัติ (Nominal Scale) ซึ่งถือเป็นระดับข้อมูลที่หยาบที่สุด

2) ข้อมูลระดับเรียงลำดับ (Ordinal Scale)

ข้อมูลระดับเรียงลำดับ เป็นระดับการวัดที่สูงกว่าข้อมูลระดับนามบัญญัติ โดยกำหนดตัวเลขหรือสัญลักษณ์เพื่อชี้ถึง **อันดับ** ของสิ่งที่ศึกษา สามารถนำมาเรียงลำดับจากมากไปหาน้อย หรือน้อยไปหามากตามเกณฑ์ เช่น ความเข้มข้น ความหนาแน่น ความแข็งแรง หรือคุณภาพผลงาน อย่างไรก็ตาม **ไม่สามารถระบุได้ว่าความแตกต่างระหว่างแต่ละลำดับมีค่าเท่าใด** และไม่อาจยืนยันได้ว่าความแตกต่างนั้นเท่ากันในทุกช่วง อีกทั้งไม่สามารถนำไปดำเนินการทางคณิตศาสตร์ (บวก ลบ คูณ หาร) ได้โดยตรง (สุจรรยา ทรัพย์ศิริโสภณ, 2555)

ตัวอย่างข้อมูลระดับเรียงลำดับ

- ระดับการศึกษา ☐ ต่ำกว่าปริญญาตรี ☐ ปริญญาตรี ☐ ปริญญาโท ☐ ปริญญาเอก
- อายุ ☐ ต่ำกว่า 20 ปี ☐ 21-30 ปี ☐ 31-40 ปี ☐ 41-50 ปี ☐ 50 ปีขึ้นไป
- รายได้ต่อเดือน ☐ ต่ำกว่า 15,000 บาท ☐ 15,000-30,000 บาท ☐ 30,001-50,000 บาท ☐ 50,000 บาทขึ้นไป

ลักษณะสำคัญ (บุญธรรม กิจปรีดาปริสุทธิ์, 2534)

1. ใช้กฎเกณฑ์เฉพาะของตัวแปรในการจัดเรียงลำดับ
2. ไม่สามารถบอกค่าความแตกต่างเชิงปริมาณระหว่างลำดับได้
3. ใช้ตัวเลขหรือสัญลักษณ์แทนลำดับโดยไม่มีความหมายเชิงปริมาณ

ค่าสถิติที่ใช้ได้ ความถี่ (Frequency) ร้อยละ (Percentage) ฐานนิยม (Mode) มัชฌิมฐาน (Median) เปอร์เซ็นไทล์ (Percentile)

สรุปได้ว่า ข้อมูลระดับเรียงลำดับเป็นข้อมูลที่บอกได้ถึงความแตกต่างและสามารถจัดลำดับได้ เช่น การเรียงนักศึกษาตามความสูงจากมากไปน้อย แม้จะบอกได้ว่าใครสูงกว่าใคร แต่ไม่สามารถระบุได้ว่าความแตกต่างระหว่างแต่ละคนมีค่าเท่าใด

3) ข้อมูลระดับช่วง หรืออันตรภาค (Interval Scale)

ข้อมูลระดับช่วงเป็นการวัดที่สามารถ **จัดลำดับได้** และ **ระบุความแตกต่างระหว่างลำดับ** ได้อย่างชัดเจน โดยความแตกต่างในแต่ละช่วง (หน่วยการวัด) มีค่าเท่ากันทุกช่วง จึงสามารถนำมาเปรียบเทียบเชิงปริมาณได้ และสามารถนำมาบวก-ลบกันได้ แต่ **ไม่สามารถนำมาคูณ-หาร หรือเปรียบเทียบเป็นจำนวนเท่าได้** เนื่องจากข้อมูลประเภทนี้ **ไม่มีศูนย์แท้ (Absolute Zero)**

ศูนย์แท้ (Absolute Zero) หมายถึง การไม่มีอยู่จริงของคุณลักษณะนั้น ๆ ขณะที่ ข้อมูลระดับช่วงมีเพียงศูนย์สมมุติ (Relative Zero) เช่น ศูนย์องศาเซลเซียส ไม่ได้หมายถึงไม่มีอุณหภูมิ แต่เป็นจุดอ้างอิงตามมาตรฐานการวัดที่กำหนด

ตัวอย่าง

- อุณหภูมิ 0 องศาเซลเซียส → ไม่ได้หมายความว่าไม่มีอุณหภูมิ
- คะแนนสอบ 0 คะแนน → ไม่ได้หมายความว่าไม่มีความรู้ แต่สะท้อนว่าผู้เรียนทำข้อสอบชุดนั้นไม่ได้

- ระดับความคิดเห็นแบบมาตราส่วนลิเคิร์ต (Likert Scale)

มาตราส่วนลิเคิร์ตเป็นมาตราส่วนการวัดเชิงปริมาณที่พัฒนาโดย Rensis Likert (1932) นักจิตวิทยาชาวอเมริกัน ใช้เพื่อวัดทัศนคติ ความคิดเห็น หรือการรับรู้ของบุคคล โดยกำหนดข้อความ (Statement) แล้วให้ผู้ตอบเลือกคำตอบตามระดับความคิดเห็น ซึ่งมักจัดเรียงจาก “เห็นด้วยอย่างยิ่ง” ไปจนถึง “ไม่เห็นด้วยอย่างยิ่ง”

ตัวอย่างการใช้ 5 ระดับ (5-point Likert Scale)

- 1 = ไม่เห็นด้วยอย่างยิ่ง
- 2 = ไม่เห็นด้วย
- 3 = ปานกลาง/ไม่แน่ใจ
- 4 = เห็นด้วย
- 5 = เห็นด้วยอย่างยิ่ง

มาตราส่วนลิเคิร์ตมักถือว่าอยู่ในระดับ Interval Scale เนื่องจากช่วงระหว่างคะแนนแต่ละระดับมีค่าเท่ากัน (เทียม) ทำให้สามารถนำไปวิเคราะห์ด้วยสถิติพารามิเตอร์ เช่น ค่าเฉลี่ย และส่วนเบี่ยงเบนมาตรฐานได้ (แม้จะมีนักวิชาการบางส่วนเห็นว่าเป็น Ordinal Scale แต่ในทางปฏิบัติทางวิจัยทางสังคมศาสตร์นิยมถือว่าเป็น Interval)

ค่าสถิติที่ใช้ได้ การแจกแจงความถี่ ร้อยละ ฐานนิยม (Mode) มัธยฐาน (Median) เปอร์เซนไทล์ (Percentile) ค่าเฉลี่ย (Mean) ส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation)

สุจรรยา ทรัพย์ศิริโสภ (2555) อธิบายว่า ข้อมูลระดับช่วงช่วยให้สามารถระบุความแตกต่างเชิงปริมาณที่ละเอียดมากขึ้น เพราะหน่วยวัดมีค่าเท่ากัน จึงสามารถบอกได้ว่าแต่ละบุคคลแตกต่างกันมากน้อยเพียงใดเมื่อเปรียบเทียบในเชิงปริมาณ เช่น คะแนนสอบของนักศึกษาทั้ง 30 คน สามารถนำมาเปรียบเทียบได้ว่าผลต่างระหว่างคะแนนเป็นเท่าใด แม้ไม่สามารถตีความว่าใครมีความรู้มากกว่ากันก็เท่า

4) ข้อมูลระดับอัตราส่วน (Ratio Scale)

ข้อมูลระดับอัตราส่วนเป็นระดับการวัดที่สมบูรณ์ที่สุดในเชิงกายภาพและเชิงคณิตศาสตร์ เพราะสามารถระบุความแตกต่างระหว่างกลุ่ม/ลำดับได้อย่างชัดเจน สามารถเปรียบเทียบปริมาณที่มาก-น้อยได้ มีช่วงห่าง (Interval) ที่เท่ากัน และที่สำคัญคือมี **ศูนย์แท้ (Absolute Zero)** ซึ่งหมายถึงการไม่มีค่าอย่างแท้จริง เช่น ศูนย์กิโลกรัม หมายถึงไม่มีน้ำหนัก ศูนย์กิโลเมตร หมายถึง ไม่มีระยะทาง หรือมีบุตร 0 คน หมายถึงไม่มีบุตรเลย เป็นต้น (สุจรรยา ทรรศิธิริสโก, 2555)

ข้อมูลระดับนี้สามารถนำมาคำนวณทางคณิตศาสตร์ได้ครบถ้วนทั้ง **บวก ลบ คูณ และหาร** รวมทั้งสามารถเปรียบเทียบเชิงอัตราส่วนได้ เช่น คนที่มีรายได้ 20,000 บาท ต่อเดือน มีรายได้มากกว่าคนที่มีรายได้ 10,000 บาทต่อเดือน “สองเท่า” อย่างแท้จริง หรือผู้ที่มียายุ 40 ปีก็มากกว่าผู้ที่มียายุ 20 ปีเป็นสองเท่าในเชิงปริมาณที่แท้จริง

ตัวอย่างข้อมูลระดับอัตราส่วน

จำนวนบุตร 0 คน หมายถึงไม่มีบุตร (ศูนย์แท้)

รายได้ต่อเดือน.....บาท

อายุ.....ปี

น้ำหนัก.....กิโลกรัม

ระยะทาง.....กิโลเมตร

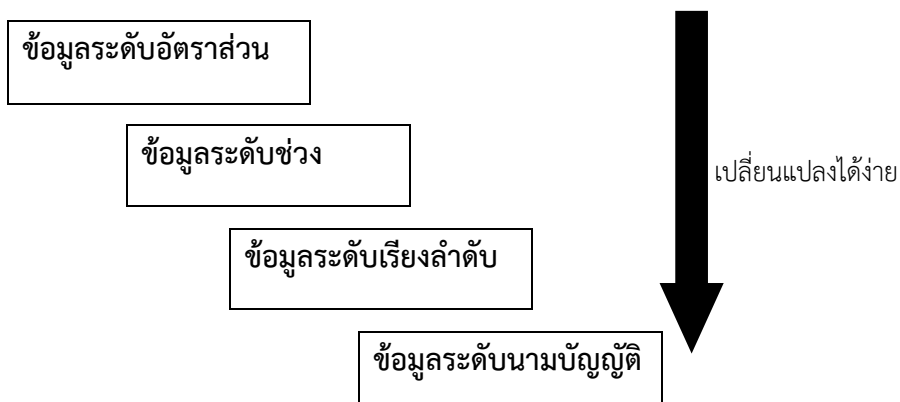
สรุปได้ว่า ข้อมูลระดับอัตราส่วนเป็นระดับข้อมูลที่สามารถจำแนกความแตกต่างได้ละเอียดที่สุด เนื่องจากมีคุณสมบัติครบทั้งการเรียงลำดับ การระบุช่วงที่เท่ากัน และมีศูนย์แท้ ทำให้สามารถนำไปใช้วิเคราะห์ทางสถิติได้ทุกวิธีการ และถือเป็นระดับการวัดที่สมบูรณ์ที่สุด (Ratio Scale)

ตารางที่ 2.1 สรุประดับของข้อมูล (Scales of Measurement)

ระดับของข้อมูล	ลักษณะสำคัญ	การดำเนินการทางคณิตศาสตร์	ค่าสถิติที่ใช้ได้	ตัวอย่าง
1. นามบัญญัติ (Nominal Scale)	ใช้จำแนกประเภทหรือกลุ่ม ไม่สามารถเรียงลำดับหรือวัดปริมาณได้ ไม่มีค่าเชิงปริมาณ	ไม่สามารถ บวก ลบ คูณ หาร ได้	ความถี่, ร้อยละ, ฐานนิยม	เพศ, ศาสนา, หมายเลขโทรศัพท์, อาชีพ
2. เรียงลำดับ (Ordinal Scale)	ใช้บอกลำดับหรืออันดับได้ แต่ไม่ทราบระยะห่างที่แน่นอนระหว่างลำดับ	ไม่สามารถบวก ลบ คูณ หาร ได้	ฐานนิยม, มัธยฐาน, เปอร์เซ็นไทล์, ความถี่	ระดับการศึกษา, ระดับรายได้, การจัดอันดับการแข่งขัน, ระดับความคิดเห็น (พอใจ-ไม่พอใจ)

ระดับของข้อมูล	ลักษณะสำคัญ	การดำเนินการทางคณิตศาสตร์	ค่าสถิติที่ใช้ได้	ตัวอย่าง
3. ช่วง/อันดับ ภาค (Interval Scale)	สามารถบอกลำดับและ ระยะห่างได้ ระยะห่าง เท่ากันทุกช่วง แต่ไม่มีศูนย์ แท้ (Relative Zero)	บวกและลบได้ แต่ไม่ สามารถคูณและหาร เชิงอัตราส่วนได้	ค่าเฉลี่ย, ส่วนเบี่ยงเบน มาตรฐาน, ความถี่, มัธยฐาน, ฐานนิยม	คะแนนสอบ, อุณหภูมิ (องศาเซลเซียส, ฟาเรน ไฮต์), มาตราส่วนลิเคิร์ท (Likert Scale)
4. อัตราส่วน (Ratio Scale)	มีคุณสมบัติครบทุกอย่าง: บอกลำดับ ระยะห่าง และ มีศูนย์แท้ (Absolute Zero)	สามารถบวก ลบ คูณ หาร และเปรียบเทียบ เชิงอัตราส่วนได้	ใช้ได้ทุกค่าสถิติ เช่น ค่าเฉลี่ย, SD, สหสัมพันธ์, การทดสอบสมมติฐาน	น้ำหนัก, ส่วนสูง, อายุ, ระยะทาง, รายได้

ข้อมูลทั้ง 4 ระดับสามารถปรับเปลี่ยนได้ โดยเฉพาะการลดระดับจากสูงไปต่ำ เช่น จากข้อมูลระดับอัตราส่วน (Ratio Scale) ไปสู่นามบัญญัติ (Nominal Scale) ทำได้ง่ายกว่า ขณะที่การยกระดับจากต่ำไปสูง เช่น จาก Ordinal ไป Interval หรือ Ratio มักซับซ้อนและไม่เหมาะสม อย่างไรก็ตาม ระดับข้อมูลที่สูงกว่าย่อมครอบคลุมคุณสมบัติของระดับที่ต่ำกว่าเสมอ ดังแสดงในแผนภาพที่ 2.1 ความสัมพันธ์ของระดับข้อมูล (สมชาย วรภิรกิจเกษมสกุล, 2554)



แผนภาพที่ 2.1 ความสัมพันธ์ของระดับข้อมูล

ตัวอย่างการปรับระดับข้อมูล เช่น การวัดความดันโลหิตเดิมเป็นข้อมูลระดับอัตราส่วน อาจจัดกลุ่มใหม่เป็นข้อมูลระดับเรียงลำดับ ได้แก่ ต่ำ-ปกติ-สูง เพื่อให้เข้าใจง่ายขึ้น ในทางกลับกัน หากผู้วิจัยต้องการวิเคราะห์ด้วยสถิติขั้นสูง ตัวแปรที่เก็บในระดับเรียงลำดับ เช่น ความคิดเห็น อาจปรับเป็นข้อมูลระดับช่วงโดยใช้มาตราส่วนลิเคิร์ท (Likert Scale) กำหนดค่า 1-5 ตั้งแต่ “ไม่เห็นด้วยอย่างยิ่ง” ถึง “เห็นด้วยอย่างยิ่ง” ทั้งนี้การเพิ่มระดับข้อมูลควรระมัดระวัง เนื่องจากอาจทำให้การตีความคลาดเคลื่อนและไม่เป็นไปตามหลักการวัด (สุทิน ชนะบุญ, 2560)

สรุปได้ว่า ระดับข้อมูลทั้ง 4 ประเภทเป็นพื้นฐานสำคัญของการวิจัยทางสังคมศาสตร์ โดยการเลือกใช้วิธีวัดและการวิเคราะห์ทางสถิติขึ้นอยู่กับระดับของข้อมูล หากข้อมูลอยู่ในระดับสูง เช่น อัตราส่วน (Ratio) จะสามารถใช้สถิติได้ครบถ้วนที่สุด ในขณะที่ข้อมูลระดับต่ำ เช่น นามบัญญัติ (Nominal) จะใช้ได้เพียงสถิติพื้นฐาน

3. การเก็บและรวบรวมข้อมูล

การเก็บและรวบรวมข้อมูลเป็นกระบวนการสำคัญที่สุดของการวิจัย เพราะข้อมูลที่ได้จะเป็นรากฐานในการวิเคราะห์และสรุปผล หากข้อมูลที่เก็บไม่มีคุณภาพหรือขาดความน่าเชื่อถือ ย่อมทำให้ข้อสรุปจากงานวิจัยขาดความน่าเชื่อถือไปด้วย ดังนั้นนักวิจัยจำเป็นต้องกำหนดขั้นตอน วิธีการ และเครื่องมือเก็บข้อมูลอย่างเป็นระบบ มีการตรวจสอบความถูกต้องและความสมบูรณ์ก่อนนำไปวิเคราะห์

ข้อมูลที่เก็บรวบรวมอาจอยู่ในรูปของข้อมูลปฐมภูมิ (Primary Data) ที่ผู้วิจัยลงมือเก็บเองโดยตรง หรือข้อมูลทุติยภูมิ (Secondary Data) ที่มีผู้อื่นเก็บไว้แล้ว นักวิจัยเพียงดึงมาใช้ให้ตรงกับวัตถุประสงค์ของการศึกษา (สุจรรยา ทรัพย์ศิริโสภา, 2555)

1) วิธีการเก็บและรวบรวมข้อมูล

(1) การเก็บรวบรวมข้อมูลจากงานทะเบียนหรือการบันทึก

เป็นการใช้ข้อมูลที่มีอยู่แล้วในหน่วยงานหรือองค์กร เช่น

- ทะเบียนราษฎร
- เวชระเบียนของโรงพยาบาล
- บัญชีการเงิน
- รายงานสถิติประจำปี

ข้อดีคือประหยัดเวลาและค่าใช้จ่าย ได้ข้อมูลต่อเนื่องในระยะยาว แต่ข้อจำกัดคือข้อมูลอาจไม่ละเอียดหรือไม่ตรงกับปัญหาที่วิจัย

(2) การเก็บรวบรวมข้อมูลจากการสำรวจ (Survey)

การสำรวจเป็นวิธีที่นิยมมากในการวิจัยทางสังคมศาสตร์และสาธารณสุข เพราะสามารถใช้กับประชากรกลุ่มใหญ่ได้ แบ่งเป็น

- การสำมะโน (Census) เก็บข้อมูลจากทุกหน่วยในประชากรที่ศึกษา ข้อดีคือครบถ้วนแม่นยำ แต่สิ้นเปลืองทรัพยากรสูง
- การสำรวจตัวอย่าง (Sample Survey) เก็บข้อมูลจากบางหน่วยของประชากรที่สุ่มมาเป็นตัวแทน ข้อดีคือประหยัดเวลาและค่าใช้จ่าย แต่ต้องออกแบบการสุ่มอย่างมีระบบเพื่อหลีกเลี่ยงความลำเอียง

- **การลงทะเบียน (Registration)** การเก็บข้อมูลอย่างต่อเนื่อง ทำให้ทราบความเปลี่ยนแปลงของประชากร เช่น ทะเบียนบ้าน การลงทะเบียนเรียน การขึ้นทะเบียนทหารกองเกิน
- **การทดลอง (Experiment)** เป็นการออกแบบการเก็บข้อมูลโดยควบคุมปัจจัยต่าง ๆ เพื่อตรวจสอบความสัมพันธ์เชิงเหตุและผล ตัวอย่างเช่น การทดลองด้านเกษตรที่ควบคุมปัจจัยเรื่องปุ๋ย น้ำ และพันธุ์พืช

(3) การเก็บรวบรวมข้อมูลจากการทดลองภาคสนามหรือโดยตรง

วิธีนี้เป็นการเก็บข้อมูลจากกลุ่มตัวอย่างหรือสิ่งที่ศึกษาโดยตรง สามารถทำได้หลายวิธี เช่น

- **การสัมภาษณ์ (Interview):** ชักถามผู้ให้ข้อมูลโดยตรง ทั้งแบบมีโครงสร้าง (Structured) และไม่มีโครงสร้าง (Unstructured)
- **การส่งแบบสอบถาม (Questionnaire/Mail):** ใช้กับกลุ่มตัวอย่างจำนวนมาก ทั้งแบบกระดาษและออนไลน์
- **การใช้โทรศัพท์ (Telephone Survey):** เหมาะกับการเก็บข้อมูลที่ต้องการความรวดเร็ว
- **การชั่ง ตวง วัด นับ (Measurement):** ใช้ในงานวิจัยเชิงปริมาณ เช่น การวัดน้ำหนัก ส่วนสูง ปริมาณผลผลิต
- **การสังเกต (Observation):** เฝ้าดูพฤติกรรมหรือปรากฏการณ์อย่างเป็นระบบ สามารถเป็นการสังเกตแบบมีส่วนร่วม หรือไม่มีส่วนร่วมได้ (กรมวิชาการเกษตร, 2558)

สรุปได้ว่า การเก็บและรวบรวมข้อมูลมีหลายรูปแบบ แต่ละวิธีมีข้อดีและข้อจำกัดต่างกัน นักวิจัยจึงต้องเลือกให้เหมาะสมกับลักษณะของปัญหาวิจัย วัตถุประสงค์ ทรัพยากร และเวลา การวางแผนและการควบคุมคุณภาพข้อมูลในขั้นตอนนี้จึงเป็นหัวใจที่จะทำให้ผลการวิเคราะห์น่าเชื่อถือและมีคุณค่าเชิงวิชาการ

2) การเก็บรวบรวมข้อมูลเชิงคุณภาพ (Qualitative Data Collection)

เป็นกระบวนการเก็บรวบรวมข้อมูลที่ไม่ใช่ตัวเลข แต่เป็นข้อมูลที่สะท้อนคุณลักษณะ ทศนคติ ความรู้สึก เหตุผล หรือแรงจูงใจของบุคคล/กลุ่มบุคคล โดยมีเป้าหมายเพื่อให้ได้ข้อมูลเชิงลึก เข้าใจปรากฏการณ์แบบองค์รวม และอธิบายบริบทอย่างละเอียด ตัวอย่างวิธีการได้แก่ การสัมภาษณ์เชิงลึก (In-depth Interview) การสนทนากลุ่ม (Focus Group Discussion) และการสังเกต (Observation)

(1) การสัมภาษณ์ (Interview)

การสัมภาษณ์เป็นวิธีการเก็บข้อมูลโดยใช้การสนทนาอย่างมีจุดประสงค์ระหว่างผู้สัมภาษณ์กับผู้ให้สัมภาษณ์ เพื่อให้ได้ความรู้ ความเข้าใจ และข้อเท็จจริงเชิงลึก สามารถสอบถามซ้ำเพื่อความชัดเจนได้ (สมชาย วรภิเษมสกุล, 2554)

(1.1) การสัมภาษณ์แบบมี/ไม่มีระบบ (Structured/Unstructured Interview)

- **แบบมีโครงสร้าง (Structured Interview)** ใช้แบบสอบถามหรือแบบสัมภาษณ์ที่สร้างไว้ล่วงหน้า มีลำดับและคำถามที่เหมือนกันทุกคน เหมาะสำหรับผู้เก็บข้อมูลที่ยังมีประสบการณ์ไม่มาก และสะดวกต่อการวิเคราะห์เชิงเปรียบเทียบ แต่ต้องระวังการจำกัดคำตอบเกินไป จึงควรมีตัวเลือกปลายเปิด

- **แบบไม่มีโครงสร้าง (Unstructured Interview)** ใช้เพียงหัวข้อหลักเป็นแนวทาง ผู้สัมภาษณ์สามารถปรับคำถามได้ตามสถานการณ์ ทำให้ได้ข้อมูลที่หลากหลายและลึกซึ้ง เหมาะกับการค้นหาความหมายที่ซ่อนเร้น แต่ต้องอาศัยทักษะและประสบการณ์สูง (นิภา ศรีไพโรจน์, 2531; Kerlinger, 1986)

(1.2) การสัมภาษณ์รายบุคคล/การสนทนากลุ่ม

- **การสัมภาษณ์รายบุคคล (Individual Interview)** เป็นการสัมภาษณ์แบบตัวต่อตัว ช่วยให้ข้อมูลเฉพาะเจาะลึก ข้อควรระวังคือผู้สัมภาษณ์ควรเป็นผู้ฟังที่ดี หลีกเลี่ยงการโต้แย้งหรือชักจูง และหลีกเลี่ยงคำถามที่ชี้นำ (ผ่องพรรณ ตรียมงคลกุล, 2543; Merriam, 1998; Van Dalen, 1979)

- **การสัมภาษณ์เป็นกลุ่มหรือการสนทนากลุ่ม (Focus Group Interview)** เป็นการสัมภาษณ์ที่ประยุกต์มาจากการอภิปรายกลุ่มผสมผสานกับวิธีการสัมภาษณ์ที่จะได้ข้อมูลเชิงปฏิสัมพันธ์ของกลุ่มบุคคลด้วย โดยมีหลักการเบื้องต้นในการดำเนินการ ดังนี้

(1) จุดมุ่งหมายของการสนทนากลุ่ม มีดังนี้

- (1.1) ใช้ในการกำหนดสมมุติฐาน
- (1.2) ใช้ในการสำรวจความคิดเห็น เจตคติและคุณลักษณะ
- (1.3) ใช้ทดสอบแนวคิดเกี่ยวกับประเด็น/ผลิตภัณฑ์ใหม่ ๆ
- (1.4) ใช้ในการประเมินผลทางธุรกิจ
- (1.5) ใช้ในการกำหนดคำถามและทดลองใช้แบบสอบถาม
- (1.6) ใช้ค้นหาคำตอบที่ยังคลุมเครือที่ได้จากเครื่องมือเก็บข้อมูลอื่น ๆ

(2) ขนาดของกลุ่มในการสัมภาษณ์ประมาณ 6-12 คน เพื่อให้สมาชิกกลุ่มได้มีโอกาสในการแสดงความคิดเห็นอย่างทั่วถึง

(3) จำนวนกลุ่ม ควรคำนึงถึงวัตถุประสงค์ของการใช้ข้อมูลว่าต้องการเปรียบเทียบหรือไม่ อาทิ ต้องการเปรียบเทียบความคิดเห็นระหว่างเพศ ควรใช้แยกกลุ่มเพื่อให้ได้ข้อมูลที่ชัดเจนและลึกซึ้ง

(4) ผู้สัมภาษณ์หรือผู้นำสนทนากลุ่ม มีบทบาทในการชี้แจงวัตถุประสงค์ และกระตุ้นให้สมาชิกกลุ่มสนทนาทุกคนแสดงความคิดเห็นเท่านั้นไม่ควรนำความคิดของตนเอง หรือเชื่อมโยงความคิดเห็นของสมาชิกกลุ่มสนทนา

(5) เวลาที่ใช้ในการสนทนากลุ่มอย่างต่อเนื่องใช้เวลาประมาณ 1- 2 ชั่วโมง ถ้าใช้เวลานานกว่านี้สมาชิกอาจจะหมดความสนใจในประเด็นที่ต้องการ

(6) การบันทึกข้อมูล ควรใช้วิธีการจดบันทึกข้อมูล ประกอบกับการบันทึกเทป/วีดิทัศน์เพื่อนำมาเก็บรายละเอียดของข้อมูล แต่ควรได้รับการอนุญาตจากสมาชิกกลุ่ม

(7) ผู้จัดบันทึกต้องบันทึกผังการนั่งสนทนา/จดบันทึกข้อมูล (เพียงอย่างเดียวเท่านั้น) และเป็นผู้ถอดเทปด้วยตนเองเพื่อให้เกิดความเข้าใจที่สอดคล้องกัน (นงพรรณ พิริยานุพงศ์, 2546)

สิน พันธุ์พินิจ (2547) ได้อธิบายกระบวนการเก็บรวบรวมข้อมูลโดยการสัมภาษณ์ ดังนี้

1) การเตรียมสัมภาษณ์ เป็นขั้นตอนในการวางแผน/กำหนดวัตถุประสงค์ว่าจะสัมภาษณ์ ใคร ที่ไหน เมื่อไร และอย่างไรให้ชัดเจน และจัดเตรียมวัสดุอุปกรณ์ ผู้ช่วยวิจัย หรือยานพาหนะ ฯลฯ

2) การสัมภาษณ์ ผู้สัมภาษณ์ควรไปให้ทันเวลากำหนดการที่ได้นัดหมายกับผู้ให้สัมภาษณ์ โดยมีการแต่งกายให้เหมาะสมกับสภาพแวดล้อมในท้องถิ่น พร้อมด้วยวัสดุอุปกรณ์ที่จัดเตรียมไว้

3) การติดตามการสัมภาษณ์ ผู้สัมภาษณ์จะต้องมีการติดตามการสัมภาษณ์ อย่างใกล้ชิดว่าได้ดำเนินการตามแผนปฏิบัติการหรือไม่ และได้จำนวนผู้ให้สัมภาษณ์ครบตามจำนวนที่ต้องการหรือไม่ (ร้อยละ 85-90) ถ้าไม่ครบตามจำนวน จะต้องมาดำเนินการตามขั้นตอนใหม่เพื่อให้ได้ข้อมูลจากผู้สัมภาษณ์ที่เพียงพอตามวัตถุประสงค์ของการวิจัย

สรุปได้ว่า การสัมภาษณ์เชิงลึก (In-depth Interview) เป็นรูปแบบเฉพาะที่นิยมในงานวิจัยคุณภาพ เพื่อค้นหาความหมาย ประสบการณ์ และมุมมองที่ซ่อนอยู่ของผู้ให้ข้อมูล การใช้เครื่องมือประกอบ เช่น การบันทึกเสียง วีดิโอ หรือการจดบันทึกภาคสนามเป็นสิ่งสำคัญเพื่อเก็บรายละเอียดและป้องกันการสูญหายของข้อมูล และความน่าเชื่อถือของข้อมูล (Credibility) ต้องอาศัยทักษะของผู้วิจัย ความเป็นกลาง และการตรวจสอบซ้ำ (Triangulation)

(2) การสังเกต (Observation)

การสังเกตเป็นวิธีการเก็บรวบรวมข้อมูลเชิงคุณภาพที่มุ่งศึกษาพฤติกรรม การกระทำ กิริยาอาการ หรือปรากฏการณ์ที่เกิดขึ้นจริง โดยใช้ประสาทสัมผัสทั้งห้า ร่วมกับเครื่องมือช่วย เช่น แบบบันทึก กล้องถ่ายรูป หรือเครื่องบันทึกเสียง เพื่อให้ได้ข้อมูลที่ถูกต้องและใกล้เคียงความจริงมากที่สุด (นงลักษณ์ วิรัชชัย, 2543)

ลักษณะสำคัญของการสังเกต คือ การเฝ้าดูอย่างมีระบบและมีวัตถุประสงค์ เพื่อวิเคราะห์ความสัมพันธ์ของสิ่งที่เกิดขึ้นกับสิ่งอื่น ๆ ในบริบทเดียวกัน (นงพรรณ พิริยานุพงศ์, 2546) ข้อควรระวังคือผู้ถูกสังเกตไม่ควรรู้ตัว เพราะอาจปรับเปลี่ยนพฤติกรรม ส่งผลให้ข้อมูลไม่สะท้อนความเป็นจริง (กฤติยา วงศ์ก้อม, 2545)

(2.1) ประเภทของการสังเกต

- จำแนกตามโครงสร้างของเครื่องมือที่ใช้

การสังเกตแบบมีโครงสร้าง (Structured Observation)

มีการกำหนดพฤติกรรมหรือปรากฏการณ์ที่ต้องการศึกษาไว้ล่วงหน้า ใช้เครื่องมือมาตรฐานเพื่อให้ข้อมูลมีความสอดคล้องและสามารถวิเคราะห์เปรียบเทียบได้ง่าย เหมาะสำหรับงานที่ต้องการข้อมูลเชิงปริมาณหรือข้อมูลที่เป็นระบบ

การสังเกตแบบไม่มีโครงสร้าง (Unstructured Observation)

ผู้สังเกตใช้วิจารณญาณเลือกประเด็นที่เห็นว่าสำคัญในสถานการณ์จริง ทำให้ได้ข้อมูลที่ยืดหยุ่นและลึกซึ้ง แต่ต้องอาศัยผู้สังเกตที่มีประสบการณ์สูงและมีความไวทางทฤษฎี (Wilkinson, 1995)

- จำแนกตามพฤติกรรมหรือบทบาทของผู้สังเกต

การสังเกตแบบมีส่วนร่วม (Participant Observation)

ผู้สังเกตเข้าไปมีบทบาทในกิจกรรมของกลุ่ม เช่น การเป็นสมาชิกโดยสมบูรณ์ (Complete Participant) หรือเข้าร่วมบางส่วน (Participant as Observer) เพื่อเข้าถึงความหมายเชิงลึกของพฤติกรรมที่เกิดขึ้น

การสังเกตแบบไม่มีส่วนร่วม (Non-Participant Observation)

ผู้สังเกตไม่เข้าไปมีส่วนร่วมในกิจกรรม เพียงเฝ้าดูพฤติกรรมจากภายนอก อาจเปิดเผยหรือซ่อนบทบาทของตนเองก็ได้ (Merriam, 1998)

- จำแนกตามบทบาทของวิธีการ

การสังเกตโดยตรง (Direct Observation) เก็บข้อมูลจาก

สถานการณ์จริงในสภาพแวดล้อมจริง

การสังเกตโดยอ้อม (Indirect Observation) เก็บข้อมูลผ่าน

บันทึก ภาพถ่าย เทป หรือสื่ออื่น ๆ

การสังเกตสถานการณ์จำลอง (Simulated Observation)
 ผู้วิจัยสร้างสถานการณ์ขึ้นเพื่อให้กลุ่มเป้าหมายแสดงพฤติกรรมที่ต้องการศึกษา
 (ปาริชาติ สถาปิตานนท์, 2546)

(2.2) ขั้นตอนในกระบวนการสังเกต

(1) **ขั้นการเตรียมการ** เป็นการเตรียมการก่อนที่จะดำเนินการสังเกต โดยกำหนดประชากรและกลุ่มตัวอย่างที่จะสังเกตตามวัตถุประสงค์ของการวิจัย จัดเตรียมแบบสังเกต และวัสดุอุปกรณ์ในการสังเกต เช่น กล้องถ่ายรูป จากนั้นประสานงานกับบุคคล หรือหน่วยงานในพื้นที่ เพื่อขออนุญาตเข้าสังเกต มีการจัดฝึกอบรมผู้สังเกตในการใช้แบบสังเกตการจดบันทึก และทักษะการอยู่ร่วมกับชุมชนและต้องจัดเตรียมค่าใช้จ่าย และยานพาหนะ และอื่น ๆ

(2) **ขั้นตอนดำเนินการ** เป็นการพูดคุย หรือแนะนำตนเองเพื่อสร้างความสัมพันธ์ที่ดีกับผู้ที่รับการสังเกตให้เกิดความเข้าใจ ไว้วางใจ และความคุ้นเคย สังเกตสภาพแวดล้อมทั่วไป และข้อมูลของกลุ่มตัวอย่างทางกายภาพในภาพรวมที่สามารถสังเกตได้ อาทิ เพศ วัย ความเป็นอยู่ และบทบาทของบุคคล เป็นต้น เมื่อเกิดความคุ้นเคยกับกลุ่มตัวอย่างแล้วจึงดำเนินการสังเกต แบบเจาะลึกกับข้อมูลที่เฉพาะเจาะจงจากแบบสังเกตที่จัดเตรียมไว้ โดยการใช้คำถาม และภาษาท้องถิ่นที่แสดงความเป็นกันเอง

(3) **ขั้นการจดบันทึกข้อมูล** เป็นการบันทึกข้อมูลที่เป็นข้อเท็จจริงที่ได้จากการสังเกตอย่างรวดเร็ว และทันทีที่พบข้อมูลเพื่อป้องกันการลืมข้อมูล หรือถ้าไม่มีโอกาสในขณะนั้นให้จดบันทึกทันทีหลังจากสิ้นสุดการดำเนินการสังเกต

(4) **ขั้นการสิ้นสุดการสังเกต** เป็นขั้นตอนหลังจากได้พิจารณาความเพียงพอของข้อมูลที่ต้องการตามวัตถุประสงค์ของการสังเกตแล้วให้แสดงการขอบคุณในความอนุเคราะห์และความร่วมมือ หรือมอบของที่ระลึก เป็นการสิ้นสุดกระบวนการสังเกต

สรุปได้ว่า การสังเกตเป็นวิธีการที่สำคัญในงานวิจัยเชิงคุณภาพ ช่วยให้ผู้วิจัยเข้าใจพฤติกรรมและปรากฏการณ์ที่เกิดขึ้นจริงอย่างเป็นธรรมชาติ แต่คุณภาพของข้อมูลขึ้นอยู่กับทักษะ ความเป็นกลาง และความละเอียดรอบคอบของผู้สังเกต จึงจำเป็นต้องมีการเตรียมการและวางแผนอย่างเป็นระบบ

ตารางที่ 2.2 เปรียบเทียบการเก็บรวบรวมข้อมูลเชิงคุณภาพ

	การสัมภาษณ์ (interview)	การสังเกต (Observation)
ข้อดี	1) แก้ปัญหาการได้รับแบบสอบถามส่งกลับคืนมาค่อนข้างน้อย 2) สามารถให้คำชี้แจงเพิ่มเติมในข้อความที่ผู้ให้ข้อมูลเกิดความสงสัยให้มีความชัดเจนมากขึ้น และซักถามข้อมูลเพิ่มเติมเมื่อพิจารณาว่ายังได้รับข้อมูลไม่ครบถ้วน 3) นำไปใช้ได้กับผู้ที่สัมภาษณ์อย่างหลากหลายลักษณะ และมีความเหมาะสมที่จะนำไปใช้กับบุคคลที่อ่านและเขียนหนังสือไม่ได้ 4) ใช้การสังเกตในระหว่างการสัมภาษณ์ทั้งบุคลิกภาพ แววตาท่าทางของผู้ให้สัมภาษณ์ และสภาวะแวดล้อมที่จะนำมาประกอบการพิจารณาสรุปผลข้อมูล 5) ผู้ให้สัมภาษณ์มีความพยายามที่จะให้ข้อมูล เนื่องจากเป็นการให้ข้อมูลแบบเผชิญหน้า	1) เป็นข้อมูลตามธรรมชาติจากแหล่งปฐมภูมิที่ได้จากพฤติกรรมการแสดงออกที่ชัดเจน ลึกซึ้ง 2) ได้รับข้อมูลที่ไม่คาดหวังหรือผู้ถูกสังเกตไม่เต็มใจในระหว่างการสังเกตที่ใช้กับผู้ที่ไม่สามารถให้ข้อมูลโดยตรง อาทิ เด็กเล็ก ๆ หรือ ผู้พิการที่พูดไม่ได้ เป็นต้น 3) สามารถแปลความหมายจากข้อมูลที่สังเกตได้ 4) เป็นหลักฐานข้อมูลเพิ่มเติมที่ใช้ในการสัมภาษณ์ให้เกิดความถูกต้องและชัดเจนมากยิ่งขึ้น
ข้อจำกัด	1) ต้องใช้เวลา แรงงาน และงบประมาณจำนวนมาก โดยเฉพาะผู้ให้สัมภาษณ์ที่อยู่ห่างไกลกัน 2) ผู้ให้สัมภาษณ์ อาจจะเกิดความล้าชวยในการให้ข้อมูลบางลักษณะที่ต้องการความเป็นส่วนตัว 3) ผู้สัมภาษณ์จะต้องเป็นบุคคลที่มีมนุษยสัมพันธ์ที่ดี จึงจะทำให้ได้รับความร่วมมือในการให้ข้อมูลที่ถูกต้อง ชัดเจน และครบถ้วน	1) บางกรณีจะใช้งบประมาณจำนวนมากหรือเวลานานกว่าวิธีเก็บข้อมูลอื่น ๆ เนื่องจากต้องรอคอยการเกิดพฤติกรรม 2) เป็นข้อมูลในเชิงบรรยายที่ยุ่งยากในการวิเคราะห์ข้อมูล เพื่อหาข้อสรุปโดยเฉพาะในเชิงเปรียบเทียบ 3) ถ้ากลุ่มตัวอย่างรู้ว่าถูกสังเกต อาจจะมีคามระมัดระวังในพฤติกรรมการแสดงออกที่ไม่สอดคล้องกับพฤติกรรมที่ต้องการ

	การสัมภาษณ์ (interview)	การสังเกต (Observation)
	4) กรณีใช้แบบสัมภาษณ์แบบไม่มีโครงสร้างทำให้ข้อมูลที่นำมาวิเคราะห์ค่อนข้างยุ่งยากโดยเฉพาะกรณีที่มีผู้สัมภาษณ์หลายคน 5) ขจัดอคติของผู้สัมภาษณ์ที่มีต่อผู้ให้สัมภาษณ์ได้ยากในกรณีรู้จักกันเป็นส่วนตัว 6) ภาษาของผู้สัมภาษณ์และผู้ให้สัมภาษณ์ในการสื่อความหมายที่ไม่สอดคล้องกันก่อให้เกิดความเข้าใจที่คลาดเคลื่อน	4) คุณภาพของข้อมูลขึ้นกับทักษะวิธีการการสังเกตของผู้สังเกตที่จะต้องได้รับการฝึกอบรม/ชี้แจงให้มีความเข้าใจที่สอดคล้องกันและมีประสิทธิภาพในการรับรู้ที่ดี 5) บางพฤติกรรมที่อาจจะเกิดขึ้นอย่างรวดเร็ว/เกิดขึ้นเพียงครั้งเดียวในขณะที่ผู้สังเกตไม่ได้สังเกต 6) ใช้กับรายบุคคลที่ไม่สามารถให้ข้อมูลด้วยวิธีการใด ๆ หรือกลุ่มตัวอย่างขนาดเล็ก 7) พฤติกรรมบางอย่างเป็นพฤติกรรมที่ซ่อนเร้น ที่จะแสดงออกนาน ๆ ครั้ง ที่ผู้สังเกตอาจจะไม่ได้รับข้อมูล

3) การเก็บรวบรวมข้อมูลเชิงปริมาณ (Quantitative Data Collection)

การรวบรวมข้อมูลเชิงปริมาณหมายถึงการเก็บข้อมูลที่อยู่ในรูปแบบ ตัวเลข (Numerical Data) ซึ่งสามารถนำไปวิเคราะห์ด้วยวิธีทางสถิติได้อย่างเป็นระบบ ข้อมูลเชิงปริมาณมีจุดเด่นคือสามารถใช้เพื่ออธิบายแนวโน้ม (Trend) วิเคราะห์ความสัมพันธ์ (Relationship) และทำนายผล (Prediction) ได้ ทำให้งานวิจัยที่อาศัยข้อมูลเชิงปริมาณมีความน่าเชื่อถือและตรวจสอบได้ (สุจรรยา ทรัพย์ศิริโสภา, 2555)

(1) แบบทดสอบ (Test)

แบบทดสอบ คือ ชุดข้อคำถามหรือสถานการณ์ที่ถูกออกแบบขึ้นเพื่อวัดพฤติกรรม ความรู้ ความสามารถ หรือคุณลักษณะของผู้ตอบตามวัตถุประสงค์ของการเรียนรู้และการวิจัย (สมชาย วรภิรมย์สกุล, 2554; บุญเชิด ภิญโญอนันตพงษ์, 2545) ผลลัพธ์ที่ได้จากแบบทดสอบจะแสดงออกมาเป็น คะแนน (Score) ซึ่งสามารถนำไปวิเคราะห์เชิงสถิติเพื่อการประเมินผลได้ โดยจำแนกประเภทของแบบทดสอบ ได้ดังนี้

(1.1) จำแนกตามวิธีสร้าง

- แบบทดสอบที่สร้างขึ้นเอง (Teacher-Made Test)

ผู้วิจัยหรือครูผู้สอนสร้างขึ้นเองเพื่อใช้วัดผลในสถานการณ์เฉพาะ เริ่มตั้งแต่การกำหนดวัตถุประสงค์ ตรวจสอบความตรงตามเนื้อหา (Content Validity) ผ่านผู้เชี่ยวชาญ และทดลองใช้เพื่อปรับปรุงคุณภาพก่อนนำไปใช้จริง

- แบบทดสอบมาตรฐาน (Standardized Test)

เป็นแบบทดสอบที่สร้างโดยผู้เชี่ยวชาญเฉพาะสาขา ผ่านกระบวนการพัฒนาอย่างเป็นระบบ มีการตรวจสอบคุณภาพจนเป็นที่ยอมรับ โดยมีเกณฑ์มาตรฐานการตรวจให้คะแนนและมีการอ้างอิงค่ามาตรฐาน (Norm) เช่น ค่ากลางของคะแนนประชากร รวมทั้งระบุค่าความเที่ยงตรง (Validity) และความเชื่อมั่น (Reliability) อย่างชัดเจน

(1.2) จำแนกตามลักษณะในการนำไปใช้

- แบบทดสอบวัดผลสัมฤทธิ์ทางการเรียน (Achievement Test)

เป็นแบบทดสอบที่ใช้วัดความสามารถด้านสติปัญญาในเนื้อหาวิชาว่า ผู้เรียนมีการเรียนรู้ที่บรรลุจุดประสงค์ที่กำหนดไว้มากหรือน้อยเพียงใด

- แบบวัดความพร้อม (Readiness Test)

เป็นแบบทดสอบที่ใช้วัดความพร้อมของผู้เรียนว่ามีความพร้อมที่จะสามารถเรียนรู้ได้ตามจุดประสงค์ที่กำหนดไว้หรือไม่ หรือมีความพร้อมใดที่ควรจะได้รับการฝึกฝนเพิ่มเติมก่อนเข้าเรียนรู้

- แบบทดสอบวินิจฉัย (Diagnostic Test)

เป็นแบบทดสอบที่ใช้วัดหรือตรวจสอบจุดบกพร่องหรือจุดเด่นของผู้เรียนในแต่ละเนื้อหาวิชา เพื่อที่จะได้นำผลการทดสอบนั้นมาใช้เพื่อแก้ไขจุดบกพร่องหรือเสริมจุดเด่นให้แก่ผู้เรียนได้อย่างมีประสิทธิภาพมากขึ้น

- แบบทดสอบเชาวน์ปัญญา (Intelligence Test)

เป็นแบบทดสอบที่ใช้วัดความสามารถด้านกระบวนการคิด หรือความสามารถในการแก้ปัญหาในสถานการณ์ใหม่ ๆ โดยใช้ประสบการณ์เดิม อาทิ แบบทดสอบเชาวน์ปัญญาของStanford-Binet หรือแบบทดสอบเชาวน์ปัญญาของWechsler เป็นต้น

- แบบทดสอบวัดความถนัด (Aptitude Test)

เป็นแบบทดสอบวัดพฤติกรรมหรือความสามารถเฉพาะด้านของผู้เรียนที่จะเกิดขึ้นในอนาคต อาทิ แบบทดสอบวัดความถนัดทางวิศวกรรม หรือแบบทดสอบวัดความถนัดทางภาษา เป็นต้น

- แบบสำรวจบุคลิกภาพ (Personality Inventories)

เป็นแบบทดสอบที่ใช้วัดคุณลักษณะ ความต้องการ การปรับตัว หรือ ค่านิยมต่าง ๆ ของบุคคล

- แบบสำรวจความสนใจด้านอาชีพ (Vocational Interest Inventories)

เป็นแบบทดสอบที่ใช้สำรวจความสนใจของบุคคลเกี่ยวกับอาชีพ หรือ งานอดิเรกที่ตนเองต้องการประกอบอาชีพหรือปฏิบัติ (สมชาย วรภิเษมสกุล, 2554)

(2) แบบสอบถาม (Questionnaire)

แบบสอบถาม เป็นเครื่องมือสำคัญในการเก็บรวบรวมข้อมูล โดยเฉพาะอย่างยิ่งในงานวิจัยทางสังคมศาสตร์ เนื่องจากสามารถใช้กับกลุ่มเป้าหมายจำนวนมากได้ในเวลาอันสั้น ผู้วิจัยไม่จำเป็นต้องพบผู้ตอบโดยตรงเสมอไป เพราะสามารถแจกจ่ายแบบสอบถามได้หลายวิธี เช่น การแจกด้วยตนเอง การส่งทางไปรษณีย์ การเก็บข้อมูลทางโทรศัพท์ หรือการใช้แบบสอบถามออนไลน์ ซึ่งสะดวก รวดเร็ว และช่วยลดค่าใช้จ่าย (ณัฐรฐนนท์ กานต์วีกุลธนา, 2565)

แบบสอบถามเป็นชุดคำถามที่ผู้วิจัยกำหนดขึ้น เพื่อใช้วัดคุณลักษณะ ความรู้ เจตคติ หรือพฤติกรรม โดยคำถามในแบบสอบถามทำหน้าที่เป็น **สิ่งเร้า (Stimulus)** ที่กระตุ้นให้ผู้ตอบสะท้อนความคิดเห็นหรือข้อมูลของตนออกมาในรูปแบบที่สามารถนำไปวิเคราะห์ได้ โดยแบ่งประเภทของแบบสอบถาม ดังนี้

(2.1) แบบสอบถามปลายเปิด (Open-Ended Form)

เป็นแบบสอบถามที่กำหนดให้เพียงข้อความเท่านั้น สำหรับคำตอบนั้นจะเป็นหน้าที่ของผู้ให้ข้อมูลที่จะได้แสดงความคิดเห็นของตนเองอย่างอิสระ และเป็นแบบสอบถามที่ตอบยากและเสียเวลาในการตอบ ซึ่งจะเห็นได้จากการสำรวจแบบสอบถามใด ๆ ที่มีคำถามปลายเปิดด้วย จะค่อนข้างไม่ได้รับคำตอบ หรือได้รับกลับคืนน้อยมาก แต่ถ้าเป็นในกรณีที่สอบถามเกี่ยวกับ แรงจูงใจ หรือสาเหตุ ก็ยังจะมีการนำมาใช้อยู่เสมอ ๆ เช่น ท่านเลือกรับราชการ เพราะเหตุใด..... อาชีพครูตามความคิดเห็นของท่านเป็นอย่างไร..... เป็นต้น (สมชาย วรภิเษมสกุล, 2554)

(2.2) แบบสอบถามปลายปิด (Close-Ended Form)

เป็นแบบสอบถามที่กำหนดทั้งคำถามและตัวเลือก โดยให้ผู้ตอบได้เลือกคำตอบจากตัวเลือกนั้น ๆ และเป็นแบบสอบถามที่ใช้เวลาในการสร้างค่อนข้างมากแต่จะสะดวกสำหรับผู้ตอบ ซึ่งข้อมูลที่ได้จะสามารถนำไปวิเคราะห์ได้ง่าย และนำเสนอได้อย่างถูกต้อง ชัดเจน โดยจำแนกตัวอย่างแบบคำถามได้ ดังนี้

(1) คำถามแบบคู่ (Dichotomous Question) มี 2 ตัวเลือก เช่น

- “คุณพอใจกับสินค้าใช้หรือไม่” (ใช่/ไม่ใช่)
- “เมืองหลวงของฝรั่งเศสคือปารีส” (จริง/เท็จ)

(2) คำถามแบบหลายตัวเลือก (Multiple Choice Question)

- คุณใช้ผลิตภัณฑ์ของเราบ่อยแค่ไหน
- ก. รายวัน
- ข. รายสัปดาห์
- ค. รายเดือน
- ง. นาน ๆ ครั้ง
- จ. ไม่เคยเลย

(3) คำถามแบบเลือกได้หลายคำตอบ (Checklist Question)

- คุณใช้แพลตฟอร์มโซเชียลมีเดียใดบ้าง (เลือกได้มากกว่า 1 ข้อ)
- ☐ Facebook
- ☐ X
- ☐ Instagram
- ☐ LinkedIn

(4) คำถามแบบมาตราส่วน (Rating Scale) ใช้วัดระดับความรู้สึก

หรือทัศนคติ ตัวอย่างที่นิยมคือ มาตราส่วนลิเคิร์ต (Likert Scale)

- ฉันพอใจกับการบริการที่ได้รับ
- ☐ 5 เห็นด้วยมากที่สุด
- ☐ 4 เห็นด้วยมาก
- ☐ 3 เห็นด้วยปานกลาง
- ☐ 2 เห็นด้วยน้อย
- ☐ 1 เห็นด้วยน้อยที่สุด

สรุปได้ว่า แบบสอบถามเป็นเครื่องมือที่มีความยืดหยุ่น สามารถเลือกใช้ได้ทั้งแบบปลายเปิดและปลายปิดตามวัตถุประสงค์ของการวิจัย โดยแบบปลายเปิดจะได้ข้อมูลเชิงลึก แต่ยากต่อการวิเคราะห์ ขณะที่แบบปลายปิดจะสะดวกต่อการวิเคราะห์เชิงสถิติและเปรียบเทียบผลได้ง่ายกว่า

ตารางที่ 2.3 เปรียบเทียบการเก็บรวบรวมข้อมูลเชิงปริมาณ

	แบบทดสอบ (Test)	แบบสอบถาม (Questionnaire)
ข้อดี	1) ความเป็นปรนัย (Objectivity) มีเกณฑ์การให้คะแนนที่ชัดเจน ทำให้ลดอคติของผู้ประเมิน 2) ความน่าเชื่อถือ (Reliability) เมื่อใช้แบบทดสอบเดิมซ้ำๆ หรือใช้แบบทดสอบที่เทียบเคียงกัน จะให้ผลลัพธ์ที่สอดคล้องกัน 3) การประยุกต์ใช้ เหมาะสำหรับการคัดเลือกพนักงาน, การประเมินผลการเรียนรู้ หรือการวิจัยเชิงจิตวิทยา	1) ความยืดหยุ่น (Flexibility) สามารถออกแบบคำถามได้หลากหลายรูปแบบ ทั้งแบบคำตอบปิดและแบบคำตอบเปิด 2) ความสะดวกในการเก็บข้อมูล สามารถส่งผ่านช่องทางออนไลน์ได้ง่าย ทำให้เข้าถึงกลุ่มตัวอย่างขนาดใหญ่ได้ในต้นทุนต่ำ 3) วัดผลได้ในวงกว้าง สามารถเก็บข้อมูลเกี่ยวกับทัศนคติ ความคิด และความรู้สึก ซึ่งแบบทดสอบไม่สามารถทำได้
ข้อเสีย	1) ไม่สามารถวัดความคิดเห็น ความรู้สึก หรือพฤติกรรมที่ซับซ้อนได้ 2) การสร้างแบบทดสอบที่มีมาตรฐานและน่าเชื่อถือต้องอาศัยผู้เชี่ยวชาญและกระบวนการที่ซับซ้อน	1) อคติในการตอบ (Response Bias) ผู้ตอบอาจไม่ตอบตามความเป็นจริงเนื่องจากแรงกดดันทางสังคม หรือความต้องการที่จะดูดีในสายตาคนอื่น 2) ไม่ได้ได้รับความร่วมมือจากผู้ตอบ 3) การตีความที่ซับซ้อน คำตอบที่ได้ อาจไม่สะท้อนถึงเจตนาที่แท้จริงของผู้ตอบ ทำให้การตีความต้องทำอย่างระมัดระวัง

4. การตรวจสอบคุณภาพของเครื่องมือที่ใช้ในการวิจัย

การสร้างและพัฒนาเครื่องมือวิจัยที่มีคุณภาพถือเป็นหัวใจสำคัญของกระบวนการวิจัย เพราะเครื่องมือที่ดีจะช่วยให้ได้ข้อมูลที่ถูกต้อง ครบถ้วน และน่าเชื่อถือ เมื่อนำไปวิเคราะห์แล้วสามารถตอบปัญหาวิจัยได้อย่างเหมาะสม หากเครื่องมือไม่มีคุณภาพ ผลการวิจัยย่อมขาดความน่าเชื่อถือไปด้วย ดังนั้นหลังจากสร้างเครื่องมือแล้ว นักวิจัยต้องดำเนินการตรวจสอบคุณภาพทั้งในด้านความเที่ยงตรง (Validity) และความเชื่อมั่น (Reliability) ผ่านขั้นตอนที่เป็นระบบ โดยใช้หลักเกณฑ์และวิธีการที่นักวิชาการยอมรับ (สมชาย วรภิเษมสกุล, 2554)

1) ความเที่ยงตรงของเครื่องมือ (Validity)

ความเที่ยงตรง หมายถึง ระดับที่เครื่องมือสามารถวัดสิ่งที่ต้องการวัดได้ตรงตามวัตถุประสงค์ของการวิจัย ซึ่งแบ่งตรวจสอบได้หลายลักษณะ ดังนี้

(1) ความเที่ยงตรงเชิงเนื้อหา (Content Validity)

เป็นการตรวจสอบว่าแบบสอบถามหรือข้อคำถามที่สร้างขึ้น ครอบคลุมเนื้อหาและวัตถุประสงค์ของการวิจัยหรือไม่ วิธีการตรวจสอบนิยมใช้ **ดัชนีความสอดคล้อง (Index of Item-Objective Congruence: IOC)** โดยนำแบบสอบถามไปให้ผู้เชี่ยวชาญอย่างน้อย 3 คนขึ้นไป พิจารณาแต่ละข้อว่ามีความเหมาะสมเพียงใด

เกณฑ์การให้คะแนนผู้เชี่ยวชาญ

+1 = ข้อคำถามสามารถใช้วัดได้

0 = ไม่แน่ใจ

-1 = ข้อคำถามไม่สามารถใช้วัดได้

สูตรการคำนวณ

$$IOC = \frac{\sum R}{N}$$

โดยที่

R = ผลการพิจารณาของผู้เชี่ยวชาญ

N = จำนวนผู้เชี่ยวชาญ

เกณฑ์การตีความ

$IOC \geq 0.50$ ขึ้นไป = ข้อคำถามมีความเหมาะสม ใช้ได้

$IOC < 0.50$ = ควรปรับปรุงหรือแก้ไขก่อนนำไปใช้

ตารางที่ 2.4 รูปแบบของแบบตรวจสอบที่ให้ผู้เชี่ยวชาญพิจารณาความเที่ยงตรงเชิงเนื้อหาของเครื่องมือที่ใช้ในการวิจัย

จุดประสงค์ที่/เนื้อหา	ข้อคำถาม	ผลการพิจารณา		
		+1	0	-1
1.	1.			
2.	2.			
	3.			
	4.			

ตัวอย่างเช่น การหาความสอดคล้องระหว่างข้อคำถามกับจุดมุ่งหมายของผู้เชี่ยวชาญ จำนวน 3 คนในการพิจารณาข้อคำถามข้อที่ 1-4 กับจุดประสงค์ข้อที่ 1 มีดังนี้

ตารางที่ 2.5 ผลการวิเคราะห์ IOC

ข้อที่	คนที่ 1			คนที่ 2			คนที่ 3			ผลรวม	$IOC = \frac{\sum R}{N}$	ผลการวิเคราะห์
	+1	0	-1	+1	0	-1	+1	0	-1			
1	✓			✓			✓			3	$\frac{3}{3} = 1$	นำไปใช้ได้
2		✓			✓			✓		0	$\frac{0}{3} = 0$	ไม่แน่ใจว่าจะวัดได้
3	✓					✓			✓	-1	$\frac{-1}{3} = -0.33$	ใช้ไม่ได้
4	✓			✓				✓		2	$\frac{2}{3} = 0.67$	นำไปใช้ได้

(2) ความเที่ยงตรงเชิงโครงสร้าง (Construct Validity)

เป็นการตรวจสอบว่าเครื่องมือสามารถวัดแนวคิดหรือโครงสร้างทางทฤษฎีที่ตั้งใจวัดได้หรือไม่ โดยใช้การวิเคราะห์ทางสถิติ เช่น การวิเคราะห์องค์ประกอบ (Factor Analysis) เพื่อดูว่าข้อคำถามแต่ละข้อสอดคล้องกับมิติของตัวแปรเชิงทฤษฎีหรือไม่

(3) ความเที่ยงตรงเชิงสภาพ (Face Validity)

เป็นการตรวจสอบเบื้องต้นจากผู้เชี่ยวชาญหรือกลุ่มตัวอย่างทดลองว่าเครื่องมือมีความชัดเจน เข้าใจง่าย ใช้ภาษาเหมาะสม และสามารถตอบได้ตรงประเด็นหรือไม่ แม้จะไม่ใช่วิธีการที่ใช้สถิติยืนยัน แต่ช่วยปรับปรุงรูปแบบและภาษาของเครื่องมือให้สมบูรณ์ขึ้น

2) การทดสอบความเชื่อถือ (Reliability)

ความเชื่อถือได้ (Reliability) หมายถึง ความคงเส้นคงวา หรือความสม่ำเสมอของผลการวัดด้วยเครื่องมือเดียวกันในหลายครั้ง ถ้าใช้เครื่องมือวัดซ้ำในกลุ่มที่มีลักษณะใกล้เคียงกันแล้วได้ผลใกล้เคียงกัน แสดงว่าเครื่องมือที่มีความน่าเชื่อถือ (สุชาติ พงศ์ศรีเจริญ, 2560)

งานวิจัยทางสังคมศาสตร์นิยมตรวจสอบความเชื่อถือของแบบสอบถามด้วยวิธี **Internal Consistency** โดยการวิเคราะห์ความสอดคล้องกันภายในของข้อคำถาม (items) ที่ใช้วัดตัวแปรเดียวกัน วิธีที่นิยมมากที่สุดคือ ค่าสัมประสิทธิ์แอลฟาของครอนบาค (Cronbach's Alpha Coefficient)

Cronbach (1970) กำหนดว่า ถ้า ค่า $\alpha \geq 0.70$ ถือว่าแบบสอบถามมีความเชื่อถือได้ในระดับยอมรับ

ค่า $\alpha \geq 0.80$ = ดี

ค่า $\alpha \geq 0.90$ = ยอดเยี่ยม (Excellent Reliability)

สูตรการคำนวณ

$$\alpha = \frac{K}{K-1} = \left(1 - \frac{\sum S_i^2}{S_t^2}\right)$$

โดยที่

α = ค่าความเชื่อถือได้ของแบบสอบถาม

K = จำนวนข้อคำถามทั้งหมด

$\sum S_i^2$ = ผลรวมของความแปรปรวนของคะแนนแต่ละข้อ

S_t^2 = ความแปรปรวนของคะแนนรวมทั้งหมด

ขั้นตอนการคำนวณ

1. หาจำนวนข้อคำถามทั้งหมด K
2. คำนวณค่าความแปรปรวน (Variance) ของแต่ละข้อ (S_i^2)
3. หาผลรวมของค่าความแปรปรวนทั้งหมด $\sum S_i^2$
4. คำนวณค่าความแปรปรวนของคะแนนรวม (S_t^2)
5. แทนค่าในสูตรเพื่อหาค่า α

ตัวอย่าง

สมมติว่านักวิจัยสร้างแบบสอบถาม 5 ข้อ เพื่อวัดความพึงพอใจของลูกค้า และเก็บข้อมูลจากลูกค้า 5 คน

ตารางที่ 2.6 ผลการทดสอบความเชื่อถือ

ผู้ตอบ	ข้อ 1	ข้อ 2	ข้อ 3	ข้อ 4	ข้อ 5	รวม (Total)
1	5	4	5	4	5	23
2	4	4	5	3	4	20
3	4	3	3	3	4	17
4	2	2	2	2	2	10
5	1	1	1	1	1	5
ความแปรปรวน	1.36	0.8	1.36	0.8	1.36	20.8

ขั้นตอนการคำนวณ

ผลรวมความแปรปรวนรายข้อ: $\sum S_i^2 = 5.68$

ความแปรปรวนรวมทั้งหมด: $S_t^2 = 20.80$

จำนวนข้อคำถาม: $K = 5$

แทนค่าในสูตร

$$\alpha = \frac{5}{4} = \left(1 - \frac{5.68}{20.80}\right)$$

$$\alpha = 1.25 \times (1 - 0.273)$$

$$\alpha = 1.25 \times 0.727 = 0.909$$

การตีความผล

- ค่าที่ได้ = 0.909
- มากกว่า 0.90 → แบบสอบถามชุดนี้มีความเชื่อถือในระดับยอดเยี่ยม
- แสดงว่าข้อคำถามทั้ง 5 ข้อมีความสอดคล้องภายใน (Internal Consistency)

สูง สามารถนำไปใช้เก็บข้อมูลจริงได้

หลักการตีความค่า Cronbach's Alpha

$\alpha \geq 0.90$ → ความเชื่อถือยอดเยี่ยม (Excellent)

$0.80 \leq \alpha < 0.90$ → ความเชื่อถือดีมาก (Good)

$0.70 \leq \alpha < 0.80$ → ความเชื่อถือใช้ได้ (Acceptable)

$0.60 \leq \alpha < 0.70$ → พอใช้ (Questionable)

$0.50 \leq \alpha < 0.60$ → ค่อนข้างต่ำ (Poor)

$\alpha < 0.50$ → ใช้ไม่ได้ (Unacceptable)

ทั้งนี้ ค่าสัมประสิทธิ์แอลฟาที่คำนวณได้นั้นจะมีค่าอยู่ระหว่าง 0 ถึง 1 ในกรณีที่ค่าสัมประสิทธิ์แอลฟา มีค่าเข้าใกล้ 1 แสดงว่าแบบสอบถามมีความเชื่อถือได้สูงหรือค่อนข้างสูง แต่ถ้าค่าสัมประสิทธิ์แอลฟา มีค่าเข้าใกล้ 0 แสดงว่าแบบสอบถามมีความเชื่อถือได้ค่อนข้างน้อย (ณัฐรฐนนท์ กานต์วิกุลธนา, 2565)

ปัจจุบัน การหาค่า Cronbach's Alpha นิยมใช้โปรแกรมสถิติ เช่น SPSS, R, JASP หรือ Python ซึ่งช่วยคำนวณได้สะดวกและแม่นยำกว่าการคำนวณมือ โดยเฉพาะเมื่อมีข้อคำถามจำนวนมาก

3) ความยาก-ง่าย และอำนาจจำแนกของข้อคำถาม

ในการสร้างและพัฒนาแบบทดสอบ (Test) นอกจากการตรวจสอบความตรง (Validity) และความเชื่อมั่น (Reliability) แล้ว ยังควรพิจารณาคุณภาพของข้อคำถามแต่ละข้อโดยใช้ดัชนีความยาก-ง่าย (Item Difficulty Index: p) และดัชนีอำนาจจำแนก (Item Discrimination Index: D) เพื่อให้มั่นใจได้ว่าข้อสอบสามารถวัดผลได้อย่างมีประสิทธิภาพ

(1) ค่าความยาก (Item Difficulty Index: p)

เป็นดัชนีที่บอกว่าข้อสอบยากหรือง่ายเพียงใด

สูตรการคำนวณ

$$p = \frac{R}{N}$$

โดยที่

R = จำนวนผู้ทำข้อสอบข้อนั้นถูก

N = จำนวนผู้เข้าสอบทั้งหมด

เกณฑ์การแปลความหมายของค่า p

$p < 0.20 \rightarrow$ ข้อยากเกินไป

$0.20 \leq p \leq 0.80 \rightarrow$ ข้อสอบมีความยากเหมาะสม

$p > 0.80 \rightarrow$ ง่ายเกินไป

ตัวอย่าง ถ้ามีผู้สอบ 100 คน และมี 65 คนตอบถูก

$$p = \frac{65}{100} = 0.65$$

ข้อสอบนี้อยู่ในช่วงที่เหมาะสม (ระดับปานกลาง)

(2) ค่าอำนาจจำแนก (Item Discrimination Index: D)

เป็นดัชนีที่บอกว่าข้อสอบสามารถแยกความแตกต่างระหว่างผู้ที่มีความสามารถสูงกับผู้ที่มีความสามารถต่ำได้ดีเพียงใด

สูตรการคำนวณ

$$D = \frac{R_u - R_l}{n}$$

โดยที่

R_u = จำนวนผู้ทำถูกในกลุ่มคะแนนสูง

R_l = จำนวนผู้ทำถูกในกลุ่มคะแนนต่ำ

n = จำนวนผู้สอบในแต่ละกลุ่ม (เช่น 27% ของผู้สอบทั้งหมด)

เกณฑ์การแปลความหมายของค่า D

$D < 0.20 \rightarrow$ ข้อสอบไม่สามารถจำแนกได้ ควรปรับปรุงหรือตัดออก

$0.20 \leq D < 0.40 \rightarrow$ อำนาจจำแนกระดับพอใช้

$0.40 \leq D < 0.70 \rightarrow$ อำนาจจำแนกระดับดี

$D \geq 0.70 \rightarrow$ อำนาจจำแนกระดับดีมาก

ตัวอย่าง ถ้าแบ่งกลุ่มสอบออกเป็น 2 กลุ่ม กลุ่มสูงและกลุ่มต่ำกลุ่มละ 30 คน

กลุ่มสูงตอบถูก 24 คน ($R_u = 24$)

กลุ่มต่ำตอบถูก 12 คน ($R_l = 12$)

$$D = \frac{24 - 12}{30} = \frac{12}{30} = 0.65$$

ข้อสอบข้อนี้มีค่า $D = 0.40$ อยู่ในเกณฑ์ที่ดี

สรุป

- ค่า p (ความยาก) ที่เหมาะสมควรอยู่ระหว่าง 0.20 – 0.80

- ค่า D (อำนาจจำแนก) ที่เหมาะสมควรมีค่า ≥ 0.20 ขึ้นไป

หากข้อสอบมีทั้งค่า p และ D อยู่ในเกณฑ์ที่ดี ข้อสอบข้อนั้นถือว่ามีคุณภาพและสามารถนำไปใช้จริงได้

สรุปได้ว่า การตรวจสอบคุณภาพของเครื่องมือวิจัยเป็นกระบวนการที่ขาดไม่ได้ เพราะช่วยยืนยันว่าเครื่องมือมีความเหมาะสม ถูกต้อง และเชื่อถือได้ การตรวจสอบประกอบด้วย

- **ความเที่ยงตรง (Validity)** เน้นตรวจสอบความสอดคล้องกับวัตถุประสงค์และทฤษฎีที่เกี่ยวข้อง

- **ความเชื่อมั่น (Reliability)** เน้นความสม่ำเสมอและความมั่นคงของผลการวัด

- **ความยาก-ง่าย และอำนาจจำแนก (สำหรับแบบทดสอบ)** ช่วยให้ข้อคำถามมีคุณภาพและสามารถวัดได้ตรงประเด็น

เครื่องมือที่ผ่านการตรวจสอบทั้ง 3 ด้านนี้ จะช่วยให้การวิจัยมีมาตรฐานสูง ผลที่ได้สามารถอ้างอิงได้ทั้งในเชิงวิชาการและการนำไปใช้เชิงปฏิบัติ

5. การสุ่มตัวอย่างและการกำหนดขนาดตัวอย่าง

ในการวิจัยเชิงสังคมศาสตร์ มนุษยศาสตร์ หรือวิทยาศาสตร์ประยุกต์ นักวิจัยมักไม่สามารถเก็บข้อมูลจาก ประชากรทั้งหมด (Population) ได้ เนื่องจากข้อจำกัดด้านเวลา งบประมาณ บุคลากร หรือความยากในการเข้าถึง ดังนั้นจึงต้องเลือกกลุ่มตัวอย่าง (Sample) ที่มีคุณสมบัติเป็น ตัวแทนของประชากร (Representative) ให้มากที่สุด เพื่อให้ผลวิจัยสามารถ สรุปอ้างอิง (Generalization) กลับไปยังประชากรได้อย่างถูกต้อง

การสุ่มตัวอย่าง (Sampling) จึงหมายถึง กระบวนการเลือกหน่วยจาก ประชากร มาเป็นตัวแทน โดยต้องใช้วิธีการที่เหมาะสมเพื่อให้ได้กลุ่มตัวอย่างที่สะท้อนลักษณะของประชากรจริง (สมชาย วรภิเษมสกุล, 2554)

1) การสุ่มตัวอย่าง

การวิจัยที่ดีจำเป็นต้องอาศัยกลุ่มตัวอย่างที่สามารถสะท้อนลักษณะของประชากรเป้าหมายได้อย่างแท้จริง เพราะการเก็บข้อมูลจากประชากรทั้งหมดมักทำได้ยาก เนื่องจากข้อจำกัดด้านเวลา งบประมาณ และทรัพยากร ดังนั้น “การสุ่มตัวอย่าง” จึงเป็นกระบวนการสำคัญที่ช่วยให้นักวิจัยได้กลุ่มตัวอย่างที่เป็นตัวแทนที่เหมาะสม เพื่อนำข้อมูลที่ได้ไปวิเคราะห์และสรุปผลอย่างมีความน่าเชื่อถือ การเลือกวิธีการสุ่มต้องพิจารณาจากลักษณะของประชากร วัตถุประสงค์การวิจัย และความเป็นไปได้เชิงปฏิบัติ โดยทั่วไป นักวิชาการแบ่งประเภทของการสุ่มกลุ่มตัวอย่างออกเป็น 2 กลุ่มใหญ่ ได้แก่ การสุ่มที่ใช้ความน่าจะเป็น (Probability Sampling) และการสุ่มที่ไม่ใช้ความน่าจะเป็น (Non-probability Sampling) ซึ่งแต่ละวิธีมีหลักการ เงื่อนไข และข้อจำกัดที่แตกต่างกันไป ดังนี้

(1) การสุ่มกลุ่มตัวอย่างที่ใช้ความน่าจะเป็น (Probability Sampling)

เป็นการสุ่มกลุ่มตัวอย่างที่สมาชิกทุก ๆ หน่วยของประชากรมีโอกาสอย่างเท่าเทียมกันที่จะเป็นตัวแทนที่ดีที่เป็นกลุ่มตัวอย่างในการวิจัย โดยข้อมูลที่รวบรวมแล้วนำมาทดสอบนัยสำคัญทางสถิติที่ใช้สถิติเชิงอ้างอิงแล้วผลการวิจัยสามารถอ้างอิงไปสู่ประชากรของการวิจัยได้ มีวิธีการสุ่ม ดังนี้ (Nachmias and, Nachmias, 1993)

1. การสุ่มกลุ่มตัวอย่างอย่างง่าย (Simple Random Sampling)

เป็นการสุ่มที่สมาชิกทุกหน่วยของประชากรที่มีจำนวนไม่มากนัก แต่มีโอกาสอย่างเท่าเทียมกัน และเป็นอิสระจากกันที่จะได้เป็นกลุ่มตัวอย่างเหมาะสมสำหรับใช้กับประชากรที่มีสภาพคล้ายคลึงกัน

2. การสุ่มกลุ่มตัวอย่างอย่างเป็นระบบ (Systematic Random Sampling)

เป็นการสุ่มตัวอย่างที่ใช้กับประชากรที่มีจำนวนมาก และรายชื่อของสมาชิกได้เรียงลำดับตามตัวอักษรหรือวิธีการที่หลากหลาย เช่น เลือกทุกๆ คนที่ 10 จากรายชื่อยกเว้นการเรียงลำดับบนพื้นฐานของค่าตัวแปรที่ศึกษาเพราะจะได้กลุ่มตัวอย่างที่แตกต่างกันอย่างชัดเจนและไม่เป็นตัวแทนที่ดีของประชากร

3. การสุ่มกลุ่มตัวอย่างแบบชั้นภูมิ (Stratified Random Sampling)

เป็นการสุ่มตัวอย่างจากประชากรที่มีจำนวนมากและมีความแตกต่างกันระหว่างหน่วยสุ่มที่สามารถจำแนกออกเป็นชั้นภูมิ (Stratum) เพื่อให้ข้อมูลที่ได้นั้นมีความครบถ้วนและครอบคลุม จะต้องดำเนินการสุ่มกลุ่มตัวอย่างจากชั้นภูมิ เช่น เพศ อายุ รายได้ แบ่งออกเป็นชั้น (Strata) แล้วสุ่มตามสัดส่วน

4. การสุ่มกลุ่มตัวอย่างแบบกลุ่ม (Cluster Random Sampling)

เป็นการสุ่มกลุ่มตัวอย่างจากประชากรที่กระจุกกระจายก่อให้เกิดความยุ่งยากในการจัดทำกรอบของประชากร หรือเป็นประชากรที่มีการรวมกลุ่มอยู่แล้วตาม

ธรรมชาติ (ตามสภาพภูมิศาสตร์/ชั้นเรียน) เช่น โรงเรียน หมู่บ้าน โดยสุ่มกลุ่มแทนการสุ่มรายบุคคล (Gall, Borg & Gall, 1996) โดยมีลักษณะในภาพรวมของแต่ละกลุ่มที่คล้ายคลึงกัน แต่ภายในกลุ่มจะมีความแตกต่างหรือความหลากหลายอย่างครบถ้วน เพื่อให้ความคลาดเคลื่อนในการประมาณค่าพารามิเตอร์ของประชากรลดลง

5. การสุ่มแบบหลายขั้นตอน (Multi-stage Sampling)

เป็นการสุ่มตัวอย่างที่มีหลายขั้นตอน มีลักษณะคล้าย ๆ กับการสุ่มตัวอย่างแบบกลุ่มที่มีหลายขั้นตอน ตั้งแต่เริ่มต้นกลุ่มใหญ่ที่สุดจนกระทั่งสิ้นสุดที่กลุ่มตัวอย่างที่ต้องการตามความเหมาะสม ดังนั้น การสุ่มแบบหลายขั้นตอนในบางครั้งนักวิชาการจึงเรียกว่าการสุ่มตัวอย่างแบบกลุ่มหลายชั้น (Multi-stage Cluster Sampling) (May, 1997) หรือเป็นการสุ่มตัวอย่างที่ใช้หลากหลายวิธีการในการสุ่มเพื่อให้ได้กลุ่มตัวอย่างที่เป็นตัวแทนของประชากรที่ซับซ้อนและมีความสอดคล้องกับความต้องการภายใต้เงื่อนไขที่จำกัด เช่น จังหวัด → อำเภอ → โรงเรียน → นักเรียน เหมาะกับประชากรขนาดใหญ่มากและซับซ้อน

(2) การสุ่มกลุ่มตัวอย่างที่ไม่ใช้ความน่าจะเป็น (Non-probability Sampling)

เป็นการสุ่มกลุ่มตัวอย่างที่ไม่ใช้หลักการของความน่าจะเป็นที่อาจจะเกิดเนื่องจากการวิจัยที่ศึกษาจากกลุ่มที่เฉพาะเจาะจงหรือมีคุณลักษณะที่สอดคล้องกับประเด็นหรือเงื่อนไขที่กำหนดไว้ หรือเนื่องจากสถานการณ์ที่แตกต่างกันไป จึงจำเป็นต้องมีการสุ่มด้วยวิธีการนี้ ในบางครั้งเรียกการสุ่มประเภทนี้ว่า “การคัดเลือก (Selection)” จำแนกได้ ดังนี้

1. วิธีการคัดเลือกแบบมีจุดประสงค์/เฉพาะเจาะจง (Purposive Selection)

เป็นการคัดเลือกกลุ่มตัวอย่างที่มีลักษณะเฉพาะเจาะจงตามหลักการของเหตุผลโดยให้ความสอดคล้องกับปัญหาการวิจัย/จุดประสงค์นั้น ๆ แต่จะต้องมีการวางแผน กำหนดขนาดกลุ่มตัวอย่าง และการเลือกกลุ่มตัวอย่างที่ดีเป็นตัวแทน ปราศจากความลำเอียง แต่ผลการวิจัยจะไม่สามารถสรุปอ้างอิงไปสู่ประชากรโดยทั่วไปได้ เช่น การศึกษาวิธีการเรียนร่วมของเด็กพิเศษกับเด็กปกติในสถานศึกษา ดังนั้นกลุ่มตัวอย่างที่นำมาศึกษาจะศึกษาเฉพาะเจาะจงในสถานศึกษาที่มีการเรียนร่วมของเด็กพิเศษกับเด็กปกติเท่านั้น เป็นต้น หรือการคัดเลือกผู้เชี่ยวชาญในการใช้เทคนิคเคลฟายที่จะต้องมีความพิถีพิถันอย่างชัดเจน มิฉะนั้นผลสรุปที่ได้อาจจะไม่น่าเชื่อถือ ฯลฯ

2. วิธีการคัดเลือกแบบกำหนดโควตา (Quota Selection)

เป็นการคัดเลือกกลุ่มตัวอย่างโดยการกำหนดสัดส่วนของจำนวนกลุ่มตัวอย่างแต่ละกลุ่มตามคุณลักษณะที่กำหนดไว้ล่วงหน้าอย่างชัดเจน แล้วเลือกตัวอย่างที่มีลักษณะดังกล่าวให้ครบตามจำนวนที่กำหนดให้เท่านั้น เช่นเดียวกับการเลือกแบบบังเอิญ

เช่น กำหนดสัดส่วนของนักศึกษาที่เป็นกลุ่มตัวอย่างให้ข้อมูลจำแนกตามชั้นปี เป็นปีที่ 1 : ปีที่ 2 : ปีที่ 3 : ปีที่ 4 ดังนี้ 35 : 30 : 20 : 15 เป็นต้น

3. วิธีการคัดเลือกแบบบังเอิญ (Accidental Selection)

เป็นการคัดเลือกกลุ่มตัวอย่างโดยบังเอิญพบหรือไม่เฉพาะเจาะจง แต่กลุ่มตัวอย่างมีลักษณะเบื้องต้นบางประการที่สอดคล้องกับลักษณะของกลุ่มตัวอย่างที่กำหนดไว้ หรือเลือกบุคคลที่อยู่ใกล้ชิด หาได้ง่ายที่สุดเป็นตัวอย่างเพื่อให้ประหยัดเวลา แรงงาน และงบประมาณ (Bailey, 1987) เช่น การสำรวจเหตุการณ์มาโรงเรียนแต่เช้าของนักศึกษาที่มาโรงเรียน 20 คนแรก เป็นต้น

4. วิธีการคัดเลือกแบบลูกโซ่ (Snowball Selection)

เป็นการคัดเลือกกลุ่มตัวอย่างที่มีคุณสมบัติที่ต้องการแล้ว โดยใช้การแนะนำของกลุ่มตัวอย่างที่ระบุกลุ่มตัวอย่างที่มีลักษณะที่ใกล้เคียงกับตนเองสำหรับเก็บรวบรวมข้อมูลได้อย่างครบถ้วนและเพียงพอจึงจะยุติการเก็บรวบรวมข้อมูล

5. วิธีการคัดเลือกแบบตามสะดวก (Convenience Selection)

เป็นการเลือกกลุ่มตัวอย่างที่หาหรือพบได้ง่าย เช่น กลุ่มตัวอย่างจากการตอบแบบสอบถามที่ลงโฆษณาในหนังสือพิมพ์/นิตยสาร เป็นต้น

6. วิธีการคัดเลือกแบบอาสาสมัคร (Voluntary Selection)

เป็นการคัดเลือกกลุ่มตัวอย่างจากสมาชิกที่อาสาเข้ามามีส่วนร่วมเป็นหน่วยตัวอย่างด้วยความเต็มใจที่มีเหตุผลแตกต่างกัน เช่น ต้องการได้รับสิ่งตอบแทน/ความเต็มใจ
สรุปวิธีการสุ่มตัวอย่างแบบใช้และไม่ใช้ความน่าจะเป็นและเงื่อนไขในการใช้ ดังนี้

ตารางที่ 2.7 การสุ่มที่ใช้ความน่าจะเป็น

วิธีสุ่ม	เงื่อนไขการใช้	จุดเด่น
การสุ่มอย่างง่าย	กลุ่มไม่เกิน 1,000 คน และประชากรเป็นเอกพันธ์	ยุติธรรมที่สุด เข้าใจง่าย
การสุ่มแบบแบ่งชั้น	ประชากรมีความแตกต่างกันตามตัวแปร เช่น เพศ อายุ	ข้อมูลครบถ้วน ครอบคลุม
การสุ่มแบบแบ่งกลุ่ม	ประชากรกระจายกว้าง และจัดเป็นกลุ่มได้ เช่น เขต ภูมิศาสตร์	ประหยัดเวลาและ ค่าใช้จ่าย
การสุ่มเป็นระบบ	มีรายชื่อประชากรครบถ้วนและเรียงลำดับ	ทำได้ง่าย เร็ว
การสุ่มหลายขั้นตอน	ประชากรขนาดใหญ่และซับซ้อน	ลดต้นทุนและสะดวกใน ภาคสนาม

ตารางที่ 2.8 การสุ่มที่ไม่ใช้ความน่าจะเป็น

วิธีสุ่ม	เงื่อนไขการใช้	ตัวอย่าง
แบบเจาะจง	ต้องการกลุ่มเฉพาะเจาะจงหรือผู้ให้ข้อมูลสำคัญ	ผู้เชี่ยวชาญ, เด็กพิเศษ
แบบโควต้า	ทราบคุณลักษณะของกลุ่มและต้องการกำหนดสัดส่วนล่วงหน้า	กำหนดสัดส่วน นศ.ปี 1-4
แบบลูกโซ่	ไม่มีข้อมูลประชากร ใช้เครือข่ายแนะนำต่อกัน	การศึกษากลุ่มเฉพาะ เช่น ผู้ติดยา
แบบบังเอิญ	เลือกจากผู้ที่พบเจอได้ง่าย	นักศึกษาที่มาโรงเรียนเข้าที่ สุด
แบบสะดวก	เลือกจากกลุ่มที่เข้าถึงง่าย	ผู้ตอบแบบสอบถามจากนิตยสาร
แบบอาสาสมัคร	เลือกจากผู้ที่มีสมัครใจเข้ามา	ผู้เข้าร่วมวิจัยด้วยความเต็มใจ

ข้อเสนอแนะในการสุ่มตัวอย่าง

1) ถ้าสมาชิกทุกหน่วยของประชากรมีลักษณะที่คล้ายคลึงกัน และไม่สามารถจัดเป็นกลุ่ม ควรใช้วิธีการสุ่มอย่างง่าย หรือถ้าสุ่มอย่างง่ายแล้วเกิดความยุ่งยากในการเก็บข้อมูล อาจใช้การสุ่มแบบเป็นระบบ แต่ถ้าไม่สามารถระบุแหล่งที่อยู่ของประชากรได้ชัดเจน ก็อาจจะใช้วิธีการสุ่มแบบบังเอิญ

2) ถ้าสมาชิกทุกหน่วยของประชากรมีลักษณะแตกต่างกันโดยสามารถจำแนกเป็นกลุ่มที่เหมือนกัน และความแตกต่างนั้นจะส่งผลต่อการวิจัย ควรเลือกใช้การสุ่มแบบแบ่งชั้น หรือ ถ้าต้องการให้จำนวนกลุ่มตัวอย่างมีสัดส่วนตามประชากรก็อาจใช้วิธีการสุ่มแบบโควต้า

3) ถ้าสมาชิกทุกหน่วยของประชากรมีลักษณะที่สามารถแบ่งเป็นกลุ่ม โดยที่แต่ละกลุ่มมีลักษณะของกลุ่มที่คล้ายคลึงกัน แต่ภายในกลุ่มมีความหลากหลาย จะใช้การสุ่มแบบแบ่งกลุ่ม

4) ในการสุ่มตัวอย่างถ้ามีข้อจำกัดไม่สามารถสุ่มได้สะดวก หรือสุ่มแล้วคิดว่าไม่สามารถเก็บข้อมูลได้ หรือไม่มีพื้นฐานเกี่ยวกับลักษณะของประชากร หรือมีความสนใจเป็นรายกรณี ควรใช้วิธีการเลือกตัวอย่างแบบเจาะจงหรือแบบบังเอิญ

5) ถ้าประชากรมีขนาดใหญ่สามารถแบ่งกลุ่มได้หลายชั้น และต้องการให้กลุ่มตัวอย่างมีการกระจายอย่างทั่วถึงควรใช้การสุ่มแบบหลายขั้นตอน (พิชิต ฤทธิ์จรูญ, 2544)

2) การกำหนดขนาดของกลุ่มตัวอย่าง

ในการวิจัยที่ต้องอาศัยการเก็บข้อมูลจากกลุ่มตัวอย่าง สิ่งสำคัญที่สุดคือกลุ่มตัวอย่างที่เลือกมาจะต้องมีความเป็นตัวแทนของประชากร (Representation) อย่าง

แท้จริง เพราะหากกลุ่มตัวอย่างไม่สะท้อนคุณลักษณะของประชากร ผลการวิเคราะห์ที่ได้ก็จะไม่สามารถนำไปสรุปอ้างอิงได้อย่างน่าเชื่อถือ ดังนั้น การกำหนดขนาดของกลุ่มตัวอย่างที่เหมาะสมจึงถือเป็นขั้นตอนที่มีความสำคัญอย่างยิ่ง เพื่อสร้างความเที่ยงตรง (Validity) และความเชื่อมั่น (Reliability) ของผลการวิจัย

(1) ปัจจัยที่ใช้พิจารณาในการกำหนดขนาดของกลุ่มตัวอย่าง

1. ขนาดของประชากร (Population Size)

ต้องพิจารณาว่าประชากรที่ศึกษาอยู่ในขอบเขตใด มีจำนวนมากหรือน้อย เพราะหากประชากรมีจำนวนน้อยก็สามารถใช้กลุ่มตัวอย่างขนาดเล็กได้ แต่ถ้าประชากรมาก กลุ่มตัวอย่างก็ต้องมีจำนวนมากขึ้นเพื่อความแม่นยำ

2. ระดับความคลาดเคลื่อนที่ยอมรับได้ (Margin of Error) และระดับความเชื่อมั่น (Confidence Level)

ทั้งสองค่ามีความสัมพันธ์กัน กล่าวคือ หากยอมให้มีความคลาดเคลื่อน 0.05 (5%) จะได้ระดับความเชื่อมั่นที่ 95% หมายถึงผลลัพธ์จากกลุ่มตัวอย่างมีโอกาส 95% ที่จะสะท้อนลักษณะของประชากรจริง

3. ข้อตกลงเบื้องต้นของสถิติที่ใช้ (Statistical Assumptions)

การเลือกใช้สถิติในการทดสอบสมมติฐานจะมีข้อกำหนดเกี่ยวกับจำนวนกลุ่มตัวอย่าง เช่น การวิเคราะห์ความแปรปรวน (ANOVA) ต้องใช้กลุ่มตัวอย่างในแต่ละกลุ่มไม่ต่ำกว่า 30 คนเพื่อให้ผลมีความน่าเชื่อถือ

4. ประเภทของการวิจัย (Types of Research)

งานวิจัยเชิงปริมาณ (Quantitative Research) มักต้องใช้กลุ่มตัวอย่างขนาดใหญ่เพื่อให้ผลลัพธ์มีความแม่นยำสูง แต่สำหรับงานวิจัยเชิงคุณภาพ (Qualitative Research) ขนาดของกลุ่มตัวอย่างอาจไม่จำเป็นต้องใหญ่ เพียงแต่ต้องให้ได้ข้อมูลที่ลึกซึ้งและเพียงพอ

5. งบประมาณและทรัพยากรที่มีอยู่ (Budget and Resources)

จำนวนกลุ่มตัวอย่างสัมพันธ์โดยตรงกับงบประมาณและกำลังคน หากมีข้อจำกัดด้านทรัพยากร นักวิจัยอาจต้องเลือกวิธีสุ่มและขนาดที่เหมาะสมที่สุดภายใต้เงื่อนไขที่มี (สมชาย วรภิเษมสกุล, 2554)

(2) วิธีการคำนวณขนาดกลุ่มตัวอย่าง

1) การกำหนดขนาดกลุ่มตัวอย่างโดยใช้สูตรคำนวณ เช่น สูตรของ Yamane (Yamane, 1973) เป็นสูตรที่นิยมใช้มากที่สุดเมื่อต้องการประมาณสัดส่วนของประชากร โดยพิจารณาจากขนาดประชากร (N) และค่าความคลาดเคลื่อนที่ยอมรับได้ (e)

สูตรของ Yamane

$$n = \frac{N}{1 + N(e^2)}$$

เมื่อ n เป็นขนาดของกลุ่มตัวอย่าง

N เป็นขนาดของประชากร

e เป็นความคลาดเคลื่อนของการประมาณค่า (0.05 หรือ 0.01)

ตัวอย่าง มีประชากรที่ต้องการศึกษาทั้งหมด 3,000 คน และกำหนดให้มีความคลาดเคลื่อน 5% ขนาดของกลุ่มตัวอย่างที่จะใช้ในการวิจัยเรื่องนี้ จะมีขนาดเท่าไร เมื่อแทนค่าในสูตร

$$n = \frac{3000}{1 + 3000(0.05^2)} = 352.94 \approx 353$$

ดังนั้น แสดงว่าจำนวนกลุ่มตัวอย่างจะเท่ากับ 353 คน

2) การกำหนดขนาดของกลุ่มตัวอย่างโดยใช้ร้อยละของประชากร

ผู้วิจัยอาจจะใช้เกณฑ์ในการพิจารณาเป็นร้อยละของประชากรที่ต้องการศึกษา ดังนี้

- ประชากรหลักร้อยละ → ใช้กลุ่มตัวอย่างประมาณ 25%
- ประชากรหลักพัน → ใช้ประมาณ 10%
- ประชากรหลักหมื่น → ใช้ประมาณ 5%
- ประชากรหลักแสน → ใช้ประมาณ 1%

3) การกำหนดขนาดของกลุ่มตัวอย่างโดยใช้ตารางสำเร็จรูป

เป็นวิธีที่ง่ายและสะดวกที่สุดสำหรับผู้วิจัยที่ไม่ต้องการคำนวณด้วยตนเอง โดยตารางเหล่านี้มักจะคำนวณมาจากสูตรทางสถิติที่ซับซ้อนเรียบร้อยแล้ว และแสดงค่าขนาดกลุ่มตัวอย่างที่เหมาะสมตามขนาดประชากรและระดับความคลาดเคลื่อนที่กำหนดไว้ล่วงหน้า เช่น ตารางของเครซี และมอร์แกน (Krejcie and Morgan, 1970) ซึ่งเป็นตารางที่ถูกอ้างอิงบ่อยที่สุด

ตารางที่ 2.9 ขนาดของกลุ่มตัวอย่างของเครชี และมอร์แกน ที่ระดับความเชื่อมั่น 95%

N	S	N	S	N	S
10	10	220	140	1200	291
15	14	230	144	1300	297
20	19	240	148	1400	302
25	24	250	152	1500	306
30	28	260	155	1600	310
35	32	270	159	1700	313
40	36	280	162	1800	317
45	40	290	165	1900	320
50	44	300	169	2000	322
55	48	320	175	2200	327
60	52	340	181	2400	331
65	56	360	186	2600	335
70	59	380	191	2800	338
75	63	400	196	3000	341
80	66	420	201	3500	346
85	70	440	205	4000	351
90	73	460	210	4500	354
95	76	480	214	5000	357
100	80	500	217	6000	361
110	86	550	226	7000	364
120	92	600	234	8000	367
130	97	650	242	9000	368
140	103	700	248	10000	370
150	108	750	254	15000	375
160	113	800	260	20000	377
170	118	850	265	30000	379
180	123	900	269	40000	380
190	127	950	274	50000	381
200	132	1000	278	75000	382

N	S	N	S	N	S
210	136	1100	285	1000000	384

โดย N คือ จำนวนประชากร / S คือ จำนวนกลุ่มตัวอย่าง

4) การกำหนดขนาดตัวอย่างโดยใช้กฎแห่งความชัดเจน (Rule of Thumb)

เป็นการกำหนดขนาดของกลุ่มตัวอย่างโดยคำนึงถึงขนาดของประชากรในลักษณะของอัตราส่วนที่คิดเป็นร้อยละ Neuman (1991) เสนออัตราส่วนเบื้องต้น เช่น

- ประชากรน้อยกว่า 1,000 คน ใช้อัตราส่วนการสุ่มกลุ่มตัวอย่างร้อยละ 30
- ประชากรเท่ากับ 10,000 คน ใช้อัตราส่วนการสุ่มกลุ่มตัวอย่างร้อยละ 10
- ประชากรเท่ากับ 150,000 คน ใช้อัตราส่วนการสุ่มกลุ่มตัวอย่างร้อยละ 1
- ประชากรมากกว่า 10,000,000 คน ใช้อัตราส่วนการสุ่มกลุ่มตัวอย่างร้อยละ 0.025

สรุปได้ว่า กลุ่มตัวอย่างที่ดีคือกลุ่มที่มีลักษณะใกล้เคียงกับประชากรมากที่สุดทั้งในด้านคุณลักษณะและขนาด โดยการสุ่มต้องปราศจากอคติ (Bias-free) และมีจำนวนเพียงพอให้สามารถวิเคราะห์เชิงสถิติได้อย่างมีความน่าเชื่อถือ แม้ว่า “การวิจัยจากประชากรทั้งหมด” จะเป็นวิธีที่ดีที่สุด แต่ในทางปฏิบัติแทบเป็นไปไม่ได้ การกำหนดกลุ่มตัวอย่างจึงเป็นเครื่องมือสำคัญที่ช่วยให้การวิจัยมีประสิทธิภาพสูงสุด

สรุปท้ายบท

ประเภทของข้อมูลและการเก็บรวบรวมข้อมูล ถือเป็นรากฐานสำคัญของการวิจัยทางสังคมศาสตร์และสถิติ เนื่องจากข้อมูลคือหัวใจของการศึกษาที่จะนำไปสู่การวิเคราะห์และสรุปผลอย่างมีคุณภาพ ข้อมูล (Data) อาจอยู่ในรูปของตัวเลข ข้อความ หรือสัญลักษณ์ โดยจำแนกได้ทั้งข้อมูลเชิงปริมาณและข้อมูลเชิงคุณภาพ ตลอดจนแบ่งตามแหล่งที่มาเป็นข้อมูลปฐมภูมิซึ่งผู้วิจัยเก็บรวบรวมเอง และข้อมูลทุติยภูมิซึ่งเป็นข้อมูลที่มีผู้เก็บไว้แล้วและสามารถนำมาใช้ซ้ำได้

ข้อมูลแต่ละประเภทมีระดับการวัดที่แตกต่างกัน ได้แก่ ระดับนามบัญญัติซึ่งใช้เพียงเพื่อการจำแนกประเภท ระดับเรียงลำดับที่สามารถจัดอันดับแต่ยังไม่บอกระยะห่าง ระดับช่วงที่สามารถบอกได้ทั้งลำดับและระยะห่างแต่ไม่มีศูนย์แท้ และระดับอัตราส่วนซึ่งสมบูรณ์ที่สุด สามารถบอกลำดับ ระยะห่าง และมีศูนย์แท้ การจำแนกระดับข้อมูลนี้มีความสำคัญเพราะเป็นตัวกำหนดว่านักวิจัยสามารถเลือกใช้สถิติประเภทใดในการวิเคราะห์

การเก็บและรวบรวมข้อมูลถือเป็นกระบวนการที่มีความสำคัญอย่างยิ่ง เนื่องจากข้อมูลที่เก็บได้จะเป็นพื้นฐานของการสรุปผลการวิจัย หากข้อมูลขาดคุณภาพ ผลการวิจัยย่อมไม่น่าเชื่อถือ การเก็บข้อมูลสามารถทำได้หลายวิธี เช่น การสัมภาษณ์ การสังเกต การใช้แบบสอบถาม และการใช้แบบทดสอบ โดยวิธีการเหล่านี้เลือกใช้แตกต่างกันไปตามลักษณะของงานวิจัย หากเป็นข้อมูลเชิงคุณภาพจะเน้นการสัมภาษณ์เชิงลึกหรือการสังเกตเพื่อเข้าถึงบริบทและความหมาย ขณะที่ข้อมูลเชิงปริมาณจะใช้แบบสอบถามหรือแบบทดสอบเพื่อให้ได้ข้อมูลที่สามารถวิเคราะห์ทางสถิติได้

เพื่อให้ได้ข้อมูลที่ถูกต้อง เครื่องมือวิจัยจำเป็นต้องผ่านการตรวจสอบคุณภาพทั้งในด้านความเที่ยงตรง (Validity) วัดได้ตรงตามวัตถุประสงค์ ความเชื่อมั่น (Reliability) ที่แสดงถึงความสม่ำเสมอของผลการวัด รวมถึงการตรวจสอบความยากง่ายและอำนาจจำแนกของข้อคำถามในกรณีที่เป็นแบบทดสอบ เพื่อให้มั่นใจว่าเครื่องมือสามารถจำแนกความแตกต่างของผู้ตอบได้อย่างมีประสิทธิภาพ

อีกประเด็นสำคัญคือการสุ่มตัวอย่างและการกำหนดขนาดตัวอย่าง เนื่องจากข้อจำกัดด้านเวลาและงบประมาณ ทำให้นักวิจัยไม่สามารถเก็บข้อมูลจากประชากรทั้งหมดได้ จึงต้องเลือกกลุ่มตัวอย่างที่เป็นตัวแทนที่ดี วิธีการสุ่มมีทั้งการสุ่มที่ใช้ความน่าจะเป็น เช่น การสุ่มอย่างง่าย การสุ่มแบบเป็นระบบ การสุ่มแบบชั้นภูมิ การสุ่มแบบกลุ่ม และการสุ่มหลายขั้นตอน รวมถึงการสุ่มที่ไม่ใช้ความน่าจะเป็น เช่น การคัดเลือกแบบเฉพาะเจาะจง การกำหนดโควตา การเลือกแบบบังเอิญ การเลือกแบบลูกโซ่ และการเลือกตามสะดวก แต่ละวิธีมีเงื่อนไขในการใช้แตกต่างกัน การเลือกวิธีที่เหมาะสมจะช่วยให้ผลการวิจัยน่าเชื่อถือมากขึ้น

การกำหนดขนาดของกลุ่มตัวอย่างก็มีความสำคัญเช่นกัน หากกลุ่มตัวอย่างมีขนาดไม่เหมาะสม ผลการวิจัยอาจขาดความเที่ยงตรงและไม่สามารถอ้างอิงไปยังประชากรได้อย่างมั่นใจ วิธีการที่นิยมใช้ ได้แก่ การคำนวณด้วยสูตรทางสถิติ เช่น สูตรของ Yamane การใช้สัดส่วนร้อยละของประชากร การใช้ตารางสำเร็จรูป เช่น ตารางของ Krejcie และ Morgan หรือการอ้างอิงตามกฎแห่งความชัดเจน (Rule of Thumb) ปัจจัยที่ต้องพิจารณาในการกำหนดขนาดตัวอย่าง ได้แก่ ขนาดของประชากร ความคลาดเคลื่อนที่ยอมรับได้ ระดับความเชื่อมั่น ประเภทของการวิจัย รวมถึงงบประมาณที่มีอยู่

โดยสรุป บทนี้ได้ชี้ให้เห็นถึงความสำคัญของการเข้าใจประเภทของข้อมูล ระดับการวัด วิธีการเก็บรวบรวมข้อมูล การตรวจสอบคุณภาพเครื่องมือ ตลอดจนการสุ่มและกำหนดขนาดกลุ่มตัวอย่างอย่างเหมาะสม ทั้งหมดนี้คือกระบวนการพื้นฐานที่จะทำให้การวิจัยมีคุณภาพ เชื่อถือได้ และสามารถนำไปใช้ประโยชน์เชิงวิชาการและเชิงปฏิบัติได้อย่างแท้จริง

บทที่ 3

การแจกแจงข้อมูลและสถิติเชิงพรรณนา (DATA DISTRIBUTION AND DESCRIPTIVE STATISTICS)

การวิจัยทางสังคมศาสตร์และการศึกษามักเกี่ยวข้องกับข้อมูลจำนวนมากและมีความซับซ้อน หากผู้วิจัยใช้ข้อมูลดิบเพียงอย่างเดียวจะยากต่อการตีความและสื่อสารผลการวิจัย ดังนั้น จึงจำเป็นต้องมีเครื่องมือทางสถิติที่ช่วยสรุปข้อมูลให้อยู่ในรูปแบบที่เข้าใจง่าย มีความเป็นระบบ และสะท้อนลักษณะสำคัญของข้อมูลอย่างชัดเจน

สถิติเชิงพรรณนา (Descriptive Statistics) เป็นกระบวนการที่มุ่งเน้นการ “พรรณนา” หรือ “อธิบาย” ข้อมูลที่เก็บมา โดยไม่เกี่ยวข้องกับการทำนายหรืออนุมาน แต่เป็นการสรุปลักษณะของข้อมูล เช่น การหาค่ากลาง การวัดการกระจาย การศึกษารูปร่างของการแจกแจง และการนำเสนอข้อมูลในรูปกราฟ สถิติเชิงพรรณนาจึงเป็นขั้นตอนแรกๆ ที่ช่วยให้นักวิจัยเข้าใจข้อมูลได้ดีขึ้น มองเห็นแนวโน้มและรูปแบบ ก่อนที่จะก้าวไปสู่การวิเคราะห์เชิงอนุมานในขั้นต่อไป

บทนี้มุ่งอธิบายเครื่องมือพื้นฐานทางสถิติเชิงพรรณนา ได้แก่ การสร้างตารางแจกแจงความถี่ การหาค่ากลาง (ค่าเฉลี่ย มัธยฐาน ฐานนิยม) การหาค่าการกระจาย (พิสัย ส่วนเบี่ยงเบนมาตรฐาน พิสัยควอไทล์) การวิเคราะห์รูปร่างของการแจกแจง (Skewness และ Kurtosis) การจำแนกข้อมูลและเปรียบเทียบกับ การแจกแจงปกติ ตลอดจนการนำเสนอข้อมูลด้วยกราฟประเภทต่าง ๆ อันได้แก่ กราฟแท่ง กราฟวงกลม ฮิสโตแกรม และกราฟเส้น เนื้อหาเหล่านี้จะช่วยให้นักวิจัยสามารถจัดการกับข้อมูลจำนวนมากให้อยู่ในรูปแบบที่สั้น กระชับ เข้าใจง่าย และพร้อมต่อการวิเคราะห์ในระดับที่ซับซ้อนมากขึ้น

1. ตารางแจกแจงความถี่ (Frequency Distribution)

การนำเสนอข้อมูลในงานวิจัยจำเป็นต้องทำให้ข้อมูลดิบที่มีจำนวนมากและกระจัดกระจายเข้าใจได้ง่าย ตารางแจกแจงความถี่ (Frequency Distribution Table) จึงถูกนำมาใช้เพื่อจัดข้อมูลให้อยู่ในรูปแบบที่กระชับ เห็นภาพรวม และสามารถวิเคราะห์เชิงสถิติต่อไปได้

1) ความหมายของตารางแจกแจงความถี่

ตารางแจกแจงความถี่ (Frequency Distribution Table) คือ ตารางที่สรุปข้อมูลจำนวนมากให้อยู่ในรูปแบบที่เป็นระเบียบ โดยแสดงว่าค่าข้อมูลแต่ละค่า หรือช่วง

ข้อมูลแต่ละช่วง ปรากฏบ่งบ่งเพียงใด (ความถี่) การใช้ตารางลักษณะนี้ช่วยให้ผู้วิจัยเข้าใจข้อมูลได้ง่ายขึ้น สามารถเห็นภาพรวมของการกระจายข้อมูล โดยไม่จำเป็นต้องอ่านค่าดิบทุกตัวอย่าง (สุวิมล ตรีภานันท์, 2546)

การแจกแจงความถี่แบ่งได้เป็น 2 ลักษณะ คือ

1. แบบไม่จัดกลุ่ม (Ungrouped Frequency Distribution): ใช้เมื่อตัวข้อมูลมีจำนวนน้อยหรือมีค่าไม่ซ้ำกันมาก ตารางจะบอกว่าค่าตัวเลขแต่ละค่าปรากฏกี่ครั้ง

2. แบบจัดกลุ่ม (Grouped Frequency Distribution): ใช้เมื่อข้อมูลมีจำนวนมาก กระจายกว้าง นิยมแบ่งข้อมูลออกเป็น “ช่วงชั้น (class interval)” เพื่อลดจำนวนรายการ และทำให้การวิเคราะห์ง่ายขึ้น (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุโข, 2537)

ตารางแจกแจงความถี่จึงถือเป็น “จุดเริ่มต้น” ของการวิเคราะห์เชิงสถิติ เพราะใช้เป็นพื้นฐานในการหาค่าสถิติต่าง ๆ เช่น ค่าเฉลี่ย มัธยฐาน ส่วนเบี่ยงเบนมาตรฐาน และการสร้างกราฟเชิงสถิติ (ล้วน สายยศ และอังคณา สายยศ, 2538)

2) การจัดชั้นข้อมูล (Class Interval)

การจัดชั้นข้อมูลหมายถึงการแบ่งข้อมูลออกเป็นช่วง ๆ ตามเกณฑ์ที่กำหนด โดยแต่ละชั้น (Class) จะมีขีดจำกัดล่าง (Lower Class Limit) และขีดจำกัดบน (Upper Class Limit) ใช้กำหนดขอบเขตของข้อมูลที่อยู่ในช่วงนั้น และอาจใช้ขอบเขตชั้น (Class Boundaries) สำหรับข้อมูลเชิงต่อเนื่อง เช่น 50–59 อาจใช้ขอบเขต 49.5–59.5 เพื่อไม่ให้มีช่องว่างระหว่างชั้น

นอกจากนี้ยังมีแนวคิดเรื่อง จุดกึ่งกลางชั้น (Midpoint หรือ Class Mark) ซึ่งคำนวณจาก (ขีดจำกัดล่าง + ขีดจำกัดบน)/2 จุดกึ่งกลางชั้นมักถูกใช้แทนค่าตัวแทนของข้อมูลในชั้น เพื่อคำนวณหาค่าเฉลี่ยหรือสถิติอื่น ๆ ได้อย่างสะดวก (นงลักษณ์ วิรัชชัย, 2543)

การกำหนดจำนวนชั้น (Number of Classes) ควรเหมาะสม หากชั้นน้อยเกินไป ข้อมูลจะหายากและเสียรายละเอียด แต่หากมากเกินไป ข้อมูลจะกระจัดกระจายและอ่านยาก โดยทั่วไปนิยมใช้ 5–15 ชั้น หรือใช้สูตร Sturges' rule:

$$k = 1 + 3.3 \log_{10} n$$

เมื่อ n คือจำนวนข้อมูลทั้งหมด (พิศิษฐ์ ตันจวนิช, 2543)

3) ขั้นตอนการสร้างตารางแจกแจงความถี่

ขั้นตอนการสร้างตารางแจกแจงความถี่มีดังนี้ (สุวิมล ตรีภานันท์, 2546; ล้วน สายยศ และอังคณา สายยศ, 2538)

1. หาค่าพิสัย (Range)

คำนวณจากข้อมูลสูงสุด – ข้อมูลต่ำสุด เช่น ข้อมูลคะแนนสูงสุด 95 และต่ำสุด 55 → พิสัย = 40

2. กำหนดจำนวนชั้น (Number of Classes)

เลือกจำนวนชั้นให้เหมาะสม เช่น 6 ชั้น หรือใช้สูตร Sturges' rule เพื่อประมาณค่าที่เหมาะสม

3. คำนวณความกว้างชั้น (Class Interval)

ความกว้างชั้น = พิสัย ÷ จำนวนชั้น เช่น พิสัย $40 \div 6 = \sim 7$

4. กำหนดขอบเขตชั้น (Class Limits/Boundaries):

เริ่มจากค่าต่ำสุด แล้วเพิ่มไปที่ละความกว้างชั้น เช่น 55–61, 62–68, ...

5. นับความถี่ (Frequency)

ตรวจสอบจำนวนข้อมูลที่ตกอยู่ในแต่ละช่วงชั้น และบันทึกลงในตาราง

6. คำนวณค่าที่เกี่ยวข้องเพิ่มเติม

เช่น ความถี่สัมพัทธ์ (Relative Frequency = ความถี่/จำนวนข้อมูลรวม) และความถี่สะสม (Cumulative Frequency) เพื่อใช้ในการวิเคราะห์และสร้างกราฟ

ตัวอย่าง คะแนนนักเรียน 20 คน

15, 17, 17, 15, 11, 14, 15, 16, 15, 11,

14, 16, 14, 12, 18, 15, 19, 11, 17, 16

ตารางที่ 3.1 ตารางแจกแจงความถี่ (Ungrouped)

คะแนน	ความถี่
11	3
12	1
14	3
15	5
16	3
17	3
18	1
19	1

รวม = 20

4) การใช้ตารางแจกแจงความถี่เพื่ออธิบายลักษณะข้อมูล

ตารางแจกแจงความถี่มีประโยชน์หลากหลาย ได้แก่

- อธิบายการกระจุกตัวของข้อมูล ทำให้เห็นว่าข้อมูลส่วนใหญ่กระจุกตัวอยู่ที่ค่าใดหรือช่วงใด เช่น คะแนนส่วนใหญ่อยู่ระหว่าง 15–17 (ล้วน สายยศ และอังคณา สายยศ, 2538)

- ใช้เป็นฐานในการคำนวณสถิติอื่น ตารางแจกแจงที่จัดเป็นชั้นสามารถใช้หาค่าเฉลี่ย มัธยฐาน และฐานนิยมจากจุดกึ่งกลางชั้นได้ (นงลักษณ์ วิรัชชัย, 2543)

- ใช้สร้างกราฟ เช่น ฮิสโตแกรม (Histogram), กราฟเส้นหลายเหลี่ยมความถี่ (Frequency Polygon) และกราฟแท่ง (Bar Chart) ซึ่งช่วยให้ตีความแนวโน้มและรูปแบบการกระจายของข้อมูลได้ง่ายขึ้น (สุวิมล ตรีภานนท์, 2546)

- ใช้เปรียบเทียบระหว่างกลุ่ม เมื่อนำตารางของหลายกลุ่มมาเปรียบเทียบ จะเห็นความแตกต่างหรือความเหมือนของการกระจายข้อมูล (พิศิษฐ์ ตัณฑวนิช, 2543)

สรุปได้ว่า ตารางแจกแจงความถี่เป็นเครื่องมือสำคัญที่ช่วยจัดระเบียบข้อมูลจำนวนมากให้เข้าใจง่าย เห็นภาพการกระจายของข้อมูลอย่างชัดเจน และเป็นพื้นฐานสำหรับการคำนวณสถิติอื่น ๆ รวมถึงการสร้างกราฟเพื่อการวิเคราะห์ที่ลึกซึ้งและมีประสิทธิภาพมากยิ่งขึ้น

2. ค่ากลาง (Measures of Central Tendency)

ค่ากลาง คือ ค่าที่ใช้แทนหรือตัวแทนของชุดข้อมูลทั้งหมด โดยบ่งบอกแนวโน้มที่ข้อมูลมีการรวมตัวอยู่รอบ ๆ จุดหนึ่ง ทำให้เข้าใจลักษณะของข้อมูลได้โดยไม่จำเป็นต้องดูค่าทุกตัว ค่ากลางจึงเป็นเครื่องมือสำคัญในการอธิบายข้อมูลเชิงพรรณนา

ความสำคัญของค่ากลาง ได้แก่ ช่วยสรุปข้อมูลจำนวนมากให้อยู่ในค่าเดียวที่เข้าใจง่าย ใช้เปรียบเทียบระหว่างกลุ่มข้อมูล เป็นพื้นฐานในการวิเคราะห์ทางสถิติอื่น ๆ เช่น การหาค่าการกระจาย หรือการทดสอบสมมติฐานทางสถิติ

ค่ากลางที่นิยมใช้มี 3 ประเภท คือ ค่าเฉลี่ย (Mean) มัธยฐาน (Median) และฐานนิยม (Mode) มีรายละเอียดดังต่อไปนี้

1) ค่าเฉลี่ย (Mean)

ค่าเฉลี่ย คือ ผลรวมของข้อมูลทั้งหมดหารด้วยจำนวนข้อมูล เป็นค่ากลางที่ใช้กันแพร่หลายที่สุดในทางสถิติ

สูตรคำนวณ

$$\bar{X} = \frac{\sum X}{N}$$

โดยที่ $\sum X$ = ผลรวมของข้อมูลทั้งหมด, N = จำนวนข้อมูล

ตัวอย่าง คะแนนนักเรียน 5 คน 12, 14, 15, 16, 18

$$\bar{X} = \frac{12 + 14 + 15 + 16 + 18}{5} = \frac{75}{5} = 15$$

ข้อดี

- ใช้ข้อมูลทุกค่า ทำให้ได้ผลลัพธ์ที่ครอบคลุม
- ใช้ในการคำนวณสถิติอื่นได้ง่าย เช่น ส่วนเบี่ยงเบนมาตรฐาน

ข้อจำกัด

- ไวต่อค่าผิดปกติ (Outliers) เช่น ถ้ามีข้อมูล 1 ค่าเบี่ยงไปมาก ค่าเฉลี่ยจะเปลี่ยนไปทันที
- ไม่เหมาะกับข้อมูลแบบจัดชั้นที่ไม่สมมาตร

2) มัธยฐาน (Median)

มัธยฐาน คือ ค่าที่อยู่ตรงกลางของข้อมูลเมื่อเรียงลำดับจากน้อยไปมาก (หรือมากไปน้อย) ครึ่งหนึ่งของข้อมูลจะมีค่าน้อยกว่ามัธยฐาน อีกครึ่งหนึ่งจะมีค่ามากกว่ามัธยฐาน มัธยฐานจึงเหมาะกับข้อมูลที่มีการกระจายไม่สมมาตร (สุวิมล ตรีภานันท์, 2546)

วิธีหาจากข้อมูลดิบ (Ungrouped Data)

1. เรียงข้อมูลจากน้อยไปมาก
2. ถ้าจำนวนข้อมูลเป็นเลขคี่ → มัธยฐานคือค่าที่อยู่ตรงกลาง
3. ถ้าจำนวนข้อมูลเป็นเลขคู่ → มัธยฐานคือค่าเฉลี่ยของสองค่ากลาง

ตัวอย่าง

- ข้อมูล 7 ค่า: 11, 12, 14, 15, 16, 17, 18

มัธยฐาน = 15 (ค่าที่ 4)

- ข้อมูล 6 ค่า: 11, 12, 14, 15, 16, 17

มัธยฐาน = $(14+15)/2 = 14.5$

วิธีหาจากข้อมูลแจกแจงความถี่ (Grouped Data)

ใช้สูตร

$$Median = L + \left(\frac{\frac{N}{2} - CF}{f} \right) \times i$$

โดยที่

L = ขอบเขตล่างของชั้นมัธยฐาน

N = จำนวนข้อมูลทั้งหมด

CF = ความถี่สะสมก่อนชั้นมัธยฐาน

f = ความถี่ของชั้นมัธยมศึกษา

i = ความกว้างของชั้น

ตัวอย่างการแทนค่า

สมมติว่ามีข้อมูล “คะแนนสอบนักเรียน 40 คน” ดังนี้

ตารางที่ 3.2 คะแนนสอบนักเรียน

ช่วงคะแนน	ความถี่ (f)	ความถี่สะสม (CF)
50–59	5	5
60–69	10	15
70–79	15	30
80–89	7	37
90–99	3	40

- $N = 40$

- $\frac{N}{2} = 20 \rightarrow$ ตำแหน่งมัธยมศึกษาอยู่ที่คนที่ 20

- ดูจากความถี่สะสม $\rightarrow$ คนที่ 20 อยู่ในช่วง 70–79 $\rightarrow$ นี่คือนักเรียนมัธยมศึกษา

- $L = 69.5$ (ขอบเขตล่างของชั้น 70–79)

- $CF = 15$ (ความถี่สะสมก่อนหน้า)

- $f = 15$ (ความถี่ของชั้นมัธยมศึกษา)

- $i = 10$ (ความกว้างชั้น)

แทนค่าในสูตร

$$Median = 69.5 + \left(\frac{20 - 15}{15} \right) \times 10$$

$$Median = 69.5 + \left(\frac{5}{15} \right) \times 10$$

$$Median = 69.5 + 3.33 = 72.83$$

ดังนั้น มัธยมศึกษา = 72.83 คะแนน

ข้อดี

- ไม่ไวต่อค่าผิดปกติ (Outliers)
- เหมาะกับข้อมูลที่มีการกระจายไม่สมมาตร

ข้อจำกัด

- ไม่ใช่ข้อมูลทุกค่าในการคำนวณ

- ไม่สะดวกต่อการนำไปใช้คำนวณสถิติอื่น ๆ

3) ฐานนิยม (Mode)

ฐานนิยม คือ ค่าของข้อมูลที่เกิดซ้ำบ่อยที่สุด เป็นค่ากลางที่ง่ายต่อการทำความเข้าใจ และใช้ได้ทั้งกับข้อมูลเชิงปริมาณและเชิงคุณภาพ

การหาฐานนิยม

- กรณีข้อมูลดิบ ดูว่าค่าใดปรากฏบ่อยที่สุด
- กรณีข้อมูลแจกแจงความถี่ฐานนิยมคือ “ชั้นที่มีความถี่สูงสุด”

ตัวอย่าง

- ข้อมูล 7 ค่า: 11, 12, 14, 15, 15, 16, 17

ฐานนิยม = 15 (เพราะเกิดซ้ำมากที่สุด)

สูตรหาฐานนิยมจากข้อมูลแจกแจงความถี่

$$Mode = L + \left(\frac{f_1 - f_0}{(2f_1 - f_0 - f_2)} \right) \times i$$

โดยที่

L = ขอบเขตล่างของชั้นที่มีความถี่สูงสุด (Modal Class)

f_1 = ความถี่ของชั้นที่มีความถี่สูงสุด

f_0 = ความถี่ของชั้นก่อนหน้า

f_2 = ความถี่ของชั้นถัดไป

i = ความกว้างของชั้น

ตัวอย่าง ตารางข้อมูลคะแนนสอบนักเรียน

ตารางที่ 3.3 ข้อมูลคะแนนสอบนักเรียน

ช่วงคะแนน	ความถี่ (f)
50-59	5
60-69	8
70-79	12
80-89	7
90-99	3

ค่าที่ใช้

$L = 69.5$ (ขอบเขตล่างของชั้น 70-79)

$f_1 = 12$ (ความถี่สูงสุด)

$f_0 = 8$ (ความถี่ชั้นก่อนหน้า)

$f_2 = 7$ (ความถี่ชั้นถัดไป)

$i = 10$ (ความกว้างชั้น)

แทนค่าในสูตร

$$Mode = 69.5 + \left(\frac{12 - 8}{(2 \times 12 - 8 - 7)} \right) \times 10$$

$$Mode = 69.5 + \left(\frac{4}{(24 - 15)} \right) \times 10$$

$$Mode = 69.5 + \left(\frac{4}{(9)} \right) \times 10$$

$$Mode = 69.5 + 4.44 = 73.94$$

ดังนั้น ฐานนิยม ≈ 73.94 คะแนน

ข้อสังเกต

- ถ้าเป็น **ข้อมูลดิบ** ไม่จำเป็นต้องใช้สูตร $\rightarrow$ แค่ว่าค่าไหนซ้ำมากที่สุด
- ถ้าเป็น **ข้อมูลแจกแจงความถี่** ค่าฐานนิยมที่คำนวณได้จะเป็น ค่าประมาณ

เชิงตัวเลขที่อยู่ภายในช่วงชั้นฐานนิยม

ข้อดี

- ใช้ได้กับข้อมูลเชิงคุณภาพ (เช่น สีที่นิยม, ยี่ห้อที่ขายดีที่สุด)
- เข้าใจง่าย เห็นได้ชัดเจนจากการสังเกตข้อมูล

ข้อจำกัด

- ข้อมูลบางชุดอาจไม่มีฐานนิยมที่ชัดเจน
- ไม่สะดวกในการนำไปใช้คำนวณต่อ

กล่าวโดยสรุป ค่ากลางเป็นเครื่องมือสำคัญในการอธิบายลักษณะของข้อมูล โดยใช้เพียงค่าตัวแทนที่สะท้อนแนวโน้มรวมของข้อมูล ซึ่งประกอบด้วยค่าเฉลี่ยที่ใช้ข้อมูลทุกค่าแต่ไวต่อค่าผิดปกติ มีฐานที่สะท้อนค่ากึ่งกลางและทนทานต่อข้อมูลเบี่ยงเบน และฐานนิยมที่บอกค่าที่เกิดซ้ำมากที่สุด ใช้ง่ายและเข้าใจได้ชัดเจน ค่ากลางทั้งสามชนิดจึงเหมาะสมต่อการใช้ในบริบทที่แตกต่างกัน และช่วยให้การวิเคราะห์ข้อมูลมีความถูกต้องและสอดคล้องกับลักษณะข้อมูลมากยิ่งขึ้น

3. ค่าการกระจาย (Measures of Dispersion)

ค่าการกระจาย (Measures of Dispersion) หมายถึง ค่าทางสถิติที่ใช้วัดระดับความแตกต่างหรือความแปรปรวนของข้อมูลเมื่อเปรียบเทียบกับค่ากลาง เช่น ค่าเฉลี่ยหรือมัธยฐาน หากข้อมูลทั้งหมดมีค่าใกล้เคียงกัน ค่าการกระจายจะมีค่าน้อย แต่หากข้อมูลมีความแตกต่างมาก ค่าการกระจายก็จะสูงขึ้น (สุวิมล ตรีภานันท์, 2546)

ความสำคัญของการวัดการกระจาย ได้แก่

- ทำให้เข้าใจ **ความน่าเชื่อถือของค่ากลาง** เพราะค่ากลางเพียงค่าเดียวอาจไม่สะท้อนภาพรวมทั้งหมดหากข้อมูลกระจายกว้าง (นงลักษณ์ วิรัชชัย, 2543)

- ช่วยเปรียบเทียบ **กลุ่มข้อมูลต่าง ๆ** ได้ เช่น กลุ่มที่มีค่าเฉลี่ยเท่ากัน แต่ระดับการกระจายต่างกัน → กลุ่มที่การกระจายน้อยถือว่ามีความสม่ำเสมอมากกว่า (พิศิษฐ ตัณฑวนิช, 2543)

- เป็นพื้นฐานของ **สถิติอนุมาน (Inferential Statistics)** หลายชนิด เช่น การทดสอบสมมติฐาน t-test, ANOVA, การหาค่าความเชื่อมั่น และการสร้างแบบจำลองทางสถิติ (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุขโข, 2537)

1) พิสัย (Range)

พิสัยคือการวัดการกระจายที่ง่ายที่สุด โดยใช้ผลต่างระหว่างค่ามากที่สุด (Maximum) และค่าน้อยที่สุด (Minimum) ในชุดข้อมูล

สูตร

$$Range = X_{max} - X_{min}$$

ตัวอย่าง คะแนนสอบสูงสุด = 95, ต่ำสุด = 55

$$Range = 95 - 55 = 40$$

ข้อดี

- คำนวณง่าย รวดเร็ว
- ใช้เป็นตัวบ่งชี้ขอบเขตของข้อมูล

ข้อจำกัด

- ใช้เพียง 2 ค่าข้อมูล (สูงสุด-ต่ำสุด) ไม่สะท้อนลักษณะโดยรวมของข้อมูล
- ไวต่อ Outliers เช่น ถ้าค่าหนึ่งเบี่ยงไปมาก พิสัยจะสูงเกินจริง (พิศิษฐ ตัณฑวนิช, 2543)

การประยุกต์

ใช้ในสถานการณ์ที่ต้องการประเมินขอบเขตข้อมูลอย่างรวดเร็ว เช่น การกำหนดเกณฑ์คะแนนสอบ หรือการประมาณช่วงอายุของกลุ่มตัวอย่าง

2) ส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation: SD)

ส่วนเบี่ยงเบนมาตรฐานเป็นค่าการกระจายที่นิยมใช้มากที่สุด แสดงถึงระดับการกระจายของข้อมูลรอบ ๆ ค่าเฉลี่ย (Mean) โดยคำนวณจากผลต่างระหว่างแต่ละค่ากับค่าเฉลี่ย แล้วยกกำลังสองเพื่อไม่ให้ค่าติดลบ

สูตร

$$s = \sqrt{\frac{\sum(X - \bar{X})^2}{n - 1}}$$

โดยที่

X = ข้อมูลแต่ละค่า

$\bar{X}$ = ค่าเฉลี่ย

n = จำนวนข้อมูล

ตัวอย่าง ข้อมูล: 10, 12, 14

$$\bar{X} = \frac{10 + 12 + 14}{3} = \frac{36}{3} = 12$$

$$s = \sqrt{\frac{(10 - 12)^2 + (12 - 12)^2 + (14 - 12)^2}{3 - 1}}$$

$$s = \sqrt{\frac{4 + 0 + 4}{2}} = \sqrt{\frac{8}{2}} = \sqrt{4} = 2$$

ข้อดี

- ใช้ข้อมูลทุกค่ามาคำนวณ จึงสะท้อนภาพรวมได้ดี
- เป็นพื้นฐานของการวิเคราะห์ทางสถิติ เช่น z-score, t-test, ANOVA

ข้อจำกัด

- คำนวณซับซ้อนกว่าพิสัย
- ไวต่อ Outliers หากมีค่าที่เบี่ยงเบนมาก ค่า SD จะสูงผิดปกติ

การประยุกต์

นิยมใช้วัดความเสี่ยงในการเงิน เช่น การประเมินความผันผวนของราคาหุ้น หรือใช้เปรียบเทียบความแตกต่างของผลสัมฤทธิ์ทางการศึกษา

3) พิสัยควอไทล์ (Interquartile Range: IQR)

พิสัยควอไทล์ คือ ค่าที่วัดการกระจายของ “50% กลาง” ของข้อมูล โดยใช้ผลต่างระหว่างควอไทล์ที่ 3 (Q3) และควอไทล์ที่ 1 (Q1)

สูตร

$$IQR = Q3 - Q1$$

ตัวอย่าง ข้อมูล 9 ค่า: 2, 4, 6, 8, 10, 12, 14, 16, 18

$$Q1 = 4, Q2 \text{ (มัธยฐาน)} = 10, Q3 = 16$$

$$IQR = 16 - 4 = 12$$

ข้อดี

- ไม่ไวต่อ Outliers มากนัก เพราะพิจารณาเฉพาะค่ากลางของข้อมูล
- แสดงการกระจายได้ชัดเจนในกรณีข้อมูลไม่ปกติ (นงลักษณ์ วิรัชชัย, 2543)

ข้อจำกัด

- ไม่ใช่ข้อมูลทุกค่ามาคำนวณ
- ต้องเรียงข้อมูลเป็นลำดับก่อนจึงจะหาควอไทล์ได้

การประยุกต์

ใช้บ่อยในงานด้านสังคมศาสตร์และการแพทย์ เช่น การรายงานค่ามัธยฐานและ IQR ของรายได้ประชากร หรือค่าทางชีวภาพ เพื่อสะท้อนข้อมูลที่มี Outliers เช่น ค่ารายได้สูงสุดของคนรวยมาก ๆ หรือค่าผิดปกติทางการแพทย์

สรุปได้ว่า การวัดการกระจายเป็นขั้นตอนสำคัญที่ทำให้การอธิบายข้อมูลสมบูรณ์ขึ้น เพราะค่ากลางเพียงอย่างเดียวไม่สามารถแสดงความแตกต่างของข้อมูลได้ครบถ้วน พิสัยเหมาะกับการใช้เบื้องต้น ส่วนเบี่ยงเบนมาตรฐานให้ข้อมูลละเอียดและใช้ต่อยอดในสถิติขั้นสูง ขณะที่พิสัยควอไทล์เหมาะกับข้อมูลที่มี Outliers เพราะสะท้อนการกระจายของค่ากลางได้ชัดเจน

4. การวิเคราะห์รูปร่างการแจกแจง (Skewness & Kurtosis)

นอกจากการหาค่ากลางและค่าการกระจายแล้ว อีกหนึ่งวิธีที่สำคัญในการทำ ความเข้าใจกับข้อมูลคือการศึกษารูปร่างของการแจกแจง การวิเคราะห์นี้ช่วยให้เห็นลักษณะการเอียง Skewness บอกทิศทางและระดับความเอียง ส่วน Kurtosis บอกความโด่ง/ความแบนและความหนาของหางของการแจกแจง ซึ่งเป็นปัจจัยที่มีผลต่อการเลือกใช้สถิติอนุมาน เนื่องจากการทดสอบสมมติฐานจำนวนมากตั้งอยู่บนสมมติฐานว่าข้อมูลมีการแจกแจงใกล้เคียงปกติ (สุวิมล ตรีภานนท์, 2546)

1) Skewness (การกระจายเอียง)

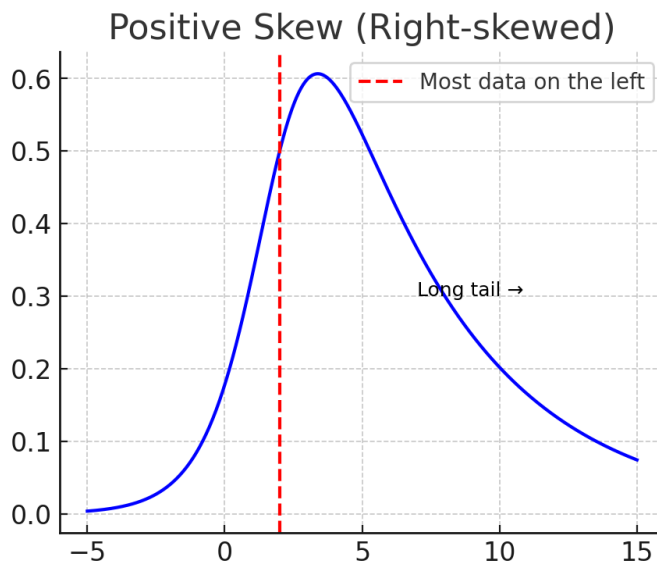
Skewness คือ สถิติที่ใช้วัดระดับความเอียงหรือความเบ้ (asymmetry) ของการแจกแจงข้อมูล เมื่อเปรียบเทียบกับ การแจกแจงปกติ (Normal Distribution) ที่มีลักษณะสมมาตร คือ ด้านซ้ายและด้านขวาของค่าเฉลี่ยมีรูปร่างใกล้เคียงกัน ถ้าการแจกแจงมีค่า Skewness เท่ากับศูนย์ แสดงว่าการแจกแจงนั้นสมมาตร ไม่เบ้ไปทางใดทางหนึ่ง (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุข, 2537)

ความสำคัญของ Skewness

- ช่วยให้นักวิจัยและนักสถิติ เข้าใจลักษณะข้อมูลเชิงลึก ว่าข้อมูลเบ้ไปทิศทางใด
- ใช้เป็น ตัวบ่งชี้ความเหมาะสมของการวิเคราะห์ เช่น หากข้อมูลเบ้มาก อาจไม่เหมาะกับการใช้สถิติพารามิเตอร์ (Parametric Statistics) ที่มีสมมติฐานว่าข้อมูลใกล้เคียงการแจกแจงปกติ
- ใช้ในการตัดสินใจว่า ควรแปลงข้อมูล (Data Transformation) หรือไม่ เช่น การใช้ Log transformation, Square root transformation เพื่อปรับข้อมูลให้ใกล้เคียงการแจกแจงปกติ

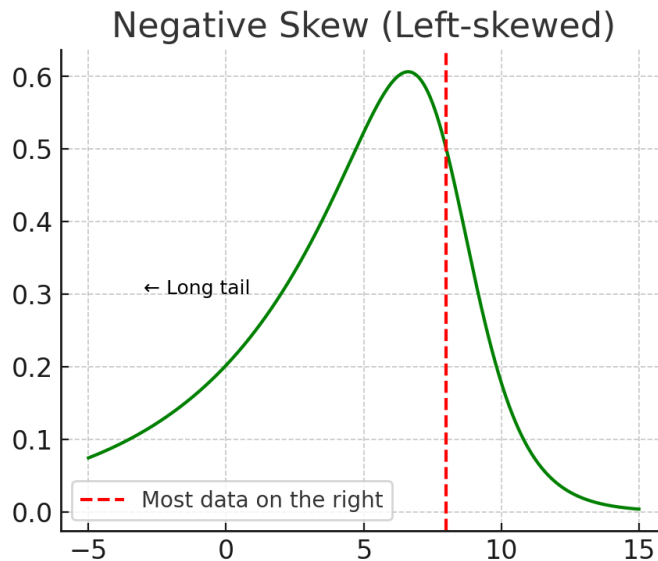
การตีความ Skewness

- ถ้า $Skewness > 0$ (Positive Skew) → การแจกแจงเบ้ไปทางขวา ข้อมูลส่วนใหญ่อยู่ทางด้านซ้าย มีหางยาวไปทางขวา เช่น การกระจายรายได้ที่คนส่วนใหญ่มีรายได้น้อย แต่มีบางคนมีรายได้สูงมาก



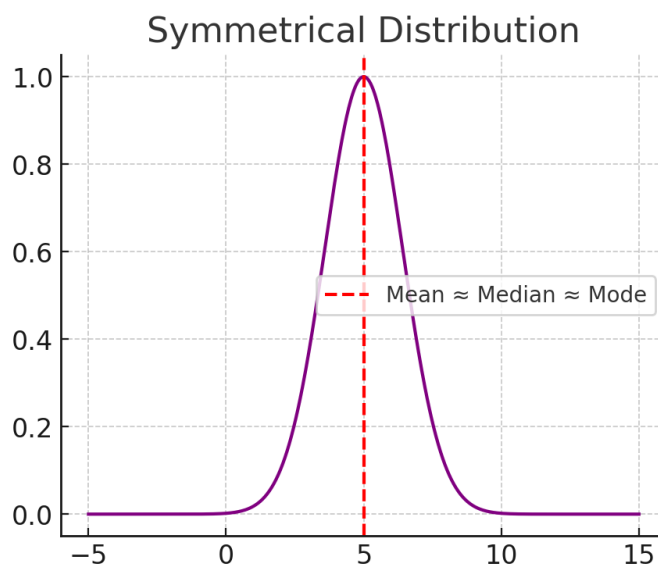
แผนภาพที่ 3.1 Positive Skew

- ถ้า $\text{Skewness} < 0$ (Negative Skew) → การแจกแจงเบ้ไปทางซ้าย
ข้อมูลส่วนใหญ่อยู่ด้านขวา มีหางยาวไปทางซ้าย เช่น คะแนนสอบที่ผู้เรียนส่วนใหญ่ทำได้สูง แต่มีบางคนสอบตก



แผนภาพที่ 3.2 Negative Skew

- ถ้า $\text{Skewness} = 0$ → การแจกแจงสมมาตร (Symmetrical) ใกล้เคียง
การแจกแจงปกติ ค่าเฉลี่ย มัธยฐาน และฐานนิยมใกล้กัน



แผนภาพที่ 3.3 Symmetrical

สูตรการคำนวณ

1. สูตรจากโมเมนต์ (Moment Skewness):

$$g_1 = \frac{m_3}{m_2^{3/2}}$$

โดยที่

$$- m_2 = \frac{1}{n} \sum (x_1 - \bar{x})^2 = \text{โมเมนต์ลำดับที่ 2 (ความแปรปรวน)}$$

$$- m_3 = \frac{1}{n} \sum (x_1 - \bar{x})^3 = \text{โมเมนต์ลำดับที่ 3}$$

$$- n = \text{จำนวนข้อมูลทั้งหมด}$$

2. สูตรสำหรับตัวอย่าง (Adjusted Skewness):

$$G_1 = \frac{\sqrt{n(n-1)}}{n-2} \times g_1 (n > 2)$$

ใช้เมื่อขนาดตัวอย่างไม่ใหญ่มาก เพื่อให้ค่ามีความเที่ยงตรงยิ่งขึ้น

ตัวอย่างการคำนวณ

ข้อมูลคะแนนสอบนักเรียน 10 คน: 12,15,17,18,20,22,25,30,35,60

$$- \text{จำนวนข้อมูล } n = 10$$

$$- \text{ค่าเฉลี่ย (Mean) } \bar{x} = 25.4$$

$$- m_2 = \frac{1}{10} \sum (x_1 - 25.4)^2 = 153.64$$

$$- m_3 = \frac{1}{10} \sum (x_1 - 25.4)^3 = 3028.44$$

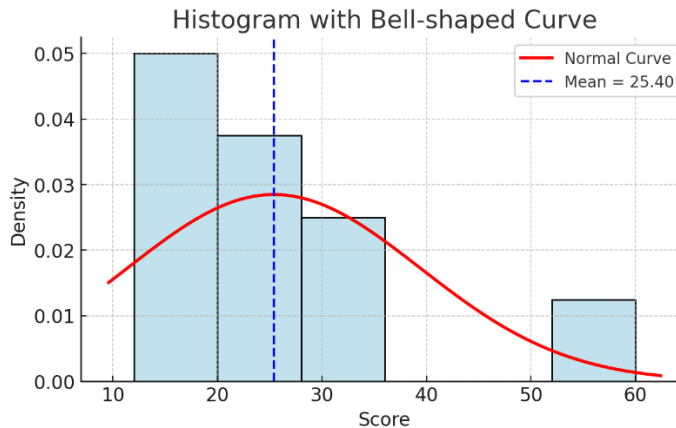
แทนค่าในสูตร

$$g_1 = \frac{3028.44}{(153.64)^{3/2}} = 1.61$$

$$G_1 = \frac{\sqrt{10 \times 9}}{8} \times 1.61 = 1.91$$

การแจกแจงนี้มีค่า Skewness > 0 แสดงว่า เบ้ขวา (Positive Skew) อย่าง

ชัดเจน



แผนภาพที่ 3.4 Histogram with Bell-shaped Curve

ด้านล่างคือ Histogram พร้อมโค้งรูประฆังคว่ำ (Bell Curve) ของข้อมูลชุดนี้

- แท่งสีน้ำเงินอ่อน = การแจกแจงข้อมูลจริง
- เส้นโค้งสีแดง = การแจกแจงปกติที่มีค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐาน

เดียวกัน

- เส้นประสีน้ำเงิน = ค่าเฉลี่ย (Mean)

การตีความ

จากกราฟ ข้อมูลส่วนใหญ่กระจุกอยู่ในช่วง 15-35 แต่มีค่า 60 ที่สูงกว่ากลุ่มมาก ทำให้หางขวายาวออกไป และค่าเฉลี่ยสูงขึ้นเมื่อเทียบกับมัธยฐาน ดังนั้น $Skewness > 0$ สะท้อนว่าเป็นการกระจายเบ้ขวา ข้อมูลมี outlier ที่ดึงการแจกแจงออกจากสมมาตร

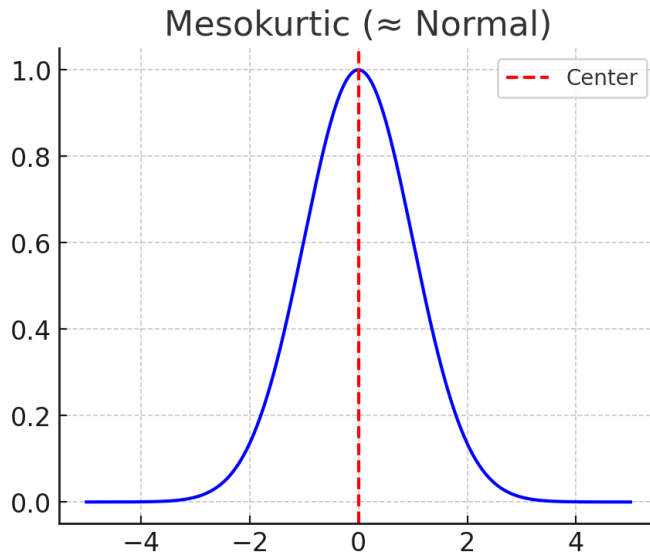
2) Kurtosis (ความโด่งของการแจกแจง)

Kurtosis (ความโด่งของการแจกแจง) คือ ค่าสถิติที่ใช้วัดลักษณะของการแจกแจงข้อมูลว่ามีความ “โด่ง” (Peakedness) หรือ “แบน” (Flatness) ของกราฟความถี่ เมื่อเปรียบเทียบกับ การแจกแจงปกติ (Normal Distribution) ที่ถือเป็นมาตรฐานอ้างอิง หากค่า **Kurtosis = 0** แสดงว่าการแจกแจงมีความสูงและการกระจายใกล้เคียงกับการแจกแจงปกติ แต่ถ้า **Kurtosis $\neq 0$** จะบ่งชี้ว่ากราฟมีลักษณะโด่งกว่าหรือแบนกว่าปกติ (พิศิษฐ์ ตัณทวนิช, 2543)

การตีความค่า Kurtosis

- Kurtosis $\approx 0 \rightarrow$ Mesokurtic

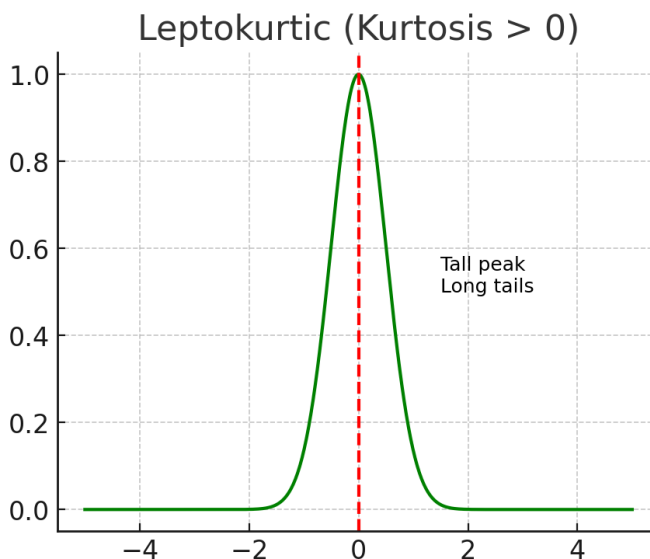
- การแจกแจงใกล้เคียงการแจกแจงปกติ
- กราฟมีความสูงสมดุล หางกระจายปกติ
- โอกาสพบ outliers อยู่ในระดับปกติ



แผนภาพที่ 3.5 Mesokurtic

- Kurtosis $> 0 \rightarrow$ Leptokurtic

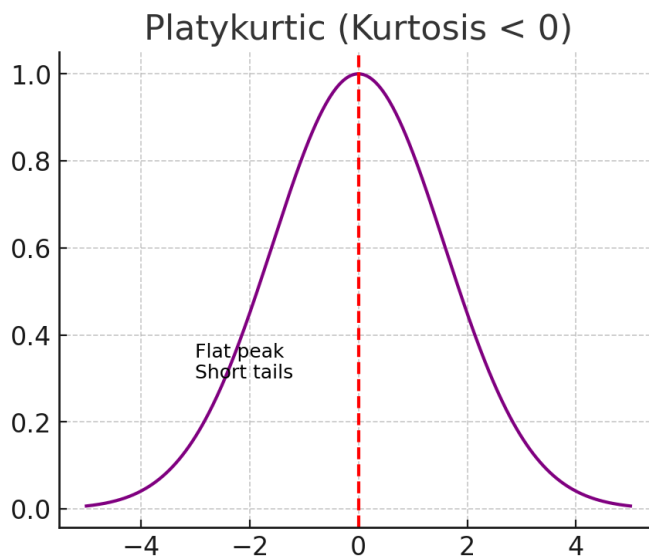
- กราฟโด่งกว่าปกติ (ยอดสูงกว่าการแจกแจงปกติ)
- ข้อมูลกระจุกอยู่รอบค่าเฉลี่ยมาก แต่มีหางยาว $\rightarrow$ แสดงว่ามีโอกาสพบข้อมูลที่แตกต่างจากกลุ่ม (outliers) ได้สูง
- ตัวอย่างการแจกแจงของผลตอบแทนการลงทุนที่มักจะมีค่ากำไร/ขาดทุนสูงผิดปกติ



แผนภาพที่ 3.6 Leptokurtic

- Kurtosis < 0 → Platykurtic

- กราฟแบนกว่าปกติ (ยอดต่ำกว่าการแจกแจงปกติ)
- ข้อมูลกระจายออกไปกว้าง หางสั้น
- แสดงว่ามีโอกาสพบ outliers น้อย เพราะค่าของข้อมูลส่วนใหญ่ไม่เบี่ยงเบนไกลจากค่าเฉลี่ยมากนัก
- ตัวอย่างคะแนนสอบที่มีการกระจายใกล้เคียงกันทุกช่วง (ไม่มีใครสูงหรือต่ำมากเกินไป)



แผนภาพที่ 3.7 Platykurtic

สูตรคำนวณ (นิยาม Fisher)

1. Excess Kurtosis (moment-based):

$$g_2 = \frac{m_4}{m_2^2} - 3$$

โดยที่

$$m_2 = \frac{1}{n} \sum (x_1 - \bar{x})^2$$

$$m_3 = \frac{1}{n} \sum (x_1 - \bar{x})^4$$

2. Adjusted Sample Excess Kurtosis:

$$G_2 = \frac{n-1}{(n-2)(n-3)} [(n+1)g_2 + 6] \quad (n > 3)$$

ตัวอย่างเชิงปฏิบัติ

1. Leptokurtic: กราฟโด่ง (Kurtosis > 0)

ข้อมูลคะแนนสอบ (14 ค่า):

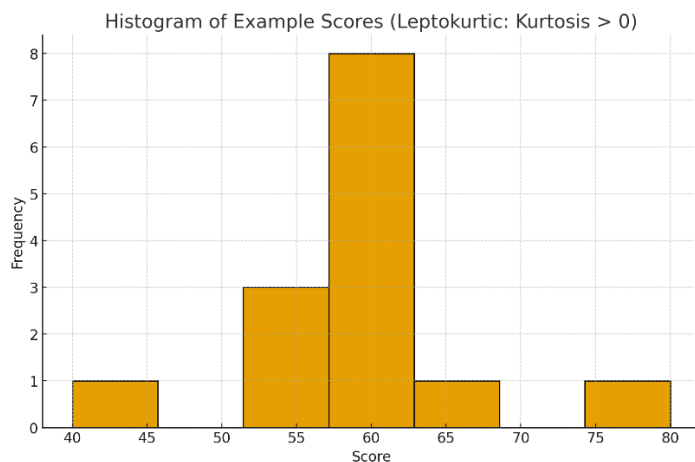
40, 55, 56, 57, 58, 60, 60, 60, 60, 60, 62, 62, 63, 80

- จำนวนข้อมูล $n = 14$
- ค่าเฉลี่ย (Mean) $\bar{x} = 59.5$
- $m_2 = 61.9643$
- $m_4 = 23002.3482$

แทนค่าในสูตร

$$g_2 = \frac{23002.3482}{(61.9643)^2} - 3 = 2.9909$$

$$G_2 = \frac{13}{(12)(11)} [(15)(2.9909) + 6] = 5.0092$$



แผนภาพที่ 3.8 Histogram of Example Scores (Leptokurtic: Kurtosis > 0)

การตีความ

กราฟโด่งกว่าปกติ ข้อมูลส่วนใหญ่กระจุกกลางใกล้ 60 แต่มีค่าปลาย (40 และ 80) ทำให้หางหนา → มีโอกาสพบ outliers มาก

2. Platykurtic: กราฟแบน (Kurtosis < 0)

ข้อมูลคะแนนสอบ (20 ค่า):

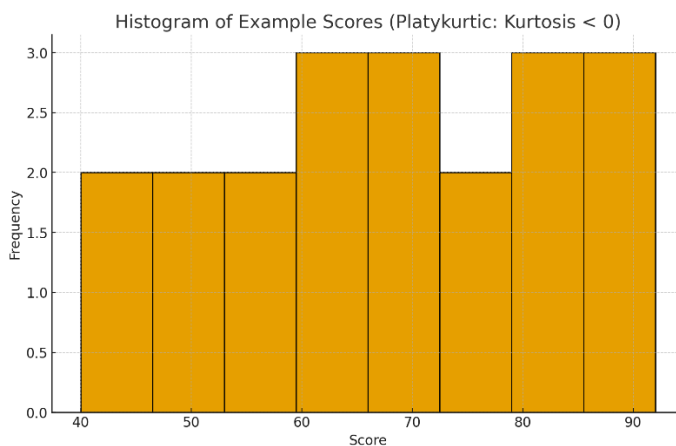
40, 45, 50, 52, 55, 58, 60, 62, 65, 68,
70, 72, 75, 78, 80, 82, 85, 88, 90, 92

- จำนวนข้อมูล $n = 20$
- ค่าเฉลี่ย (Mean) $\bar{x} = 68.35$
- $m_2 = 226.3275$
- $m_4 = 100018.3357$

แทนค่าในสูตร

$$g_2 = \frac{100018.3357}{(226.3275)^2} - 3 = -1.0474$$

$$G_2 = -0.9932$$



แผนภาพที่ 3.9 Histogram of Example Scores (Platykurtic: Kurtosis < 0)

การตีความ

กราฟแบนกว่าปกติ ข้อมูลกระจายค่อนข้างกว้างทั่วทั้งช่วง 40-92 ไม่มีการกระจุกกลางชัดเจน และหางบาง → โอกาสพบ outliers ต่ำกว่าปกติ

ตารางที่ 3.4 สรุปเปรียบเทียบประเภท Kurtosis

ประเภท	ค่า Kurtosis	ลักษณะกราฟ	ความหมาย
Mesokurtic	≈ 0	ใกล้เคียงปกติ	ข้อมูลกระจายสมมาตรตามปกติ
Leptokurtic	> 0	โด่ง หางหนา	กระจุกกลางสูง พบ outliers ง่าย
Platykurtic	< 0	แบน หางบาง	ข้อมูลกระจายกว้าง โอกาส outliers น้อย

ข้อสังเกต

- ค่า Kurtosis ช่วยบอกความเสี่ยง ในการเกิดค่าที่เบี่ยงเบนสุดขีด เช่น การเงิน หุ่น หรือการวัดผลการศึกษา

- Leptokurtic เหมาะกับการวิเคราะห์ที่ต้องจับตา outliers
- Platykurtic เชื่อว่าข้อมูล “ปลอดภัย” กว่าตามธรรมชาติของการแจกแจง
- ควรใช้ คู่กับค่า Skewness เพื่อทำความเข้าใจการแจกแจงข้อมูลให้สมบูรณ์

สรุปได้ว่า การวิเคราะห์รูปร่างการแจกแจงด้วย Skewness และ Kurtosis เป็นขั้นตอนสำคัญที่ช่วยให้นักวิจัยเข้าใจลักษณะข้อมูลได้อย่างลึกซึ้งกว่าเพียงการหาค่ากลางหรือค่าการกระจายทั่วไป โดย Skewness ใช้วัดความเอียงของข้อมูลว่ามีการเบ้ไปด้านซ้ายหรือขวา ขณะที่ Kurtosis ใช้วัดความโด่งหรือความแบนของการแจกแจง รวมถึงความหนาของหางที่สะท้อนโอกาสเกิดค่าผิดปกติ การทำความเข้าใจทั้งสองตัวชี้วัดนี้ช่วยให้เลือกใช้สถิติอนุมานได้เหมาะสม ความสำเร็จได้ตรงกับสภาพจริง และสามารถนำไปประยุกต์ใช้ในงานวิจัยด้านการศึกษา เศรษฐศาสตร์ การเงิน และสังคมศาสตร์ได้อย่างมีประสิทธิภาพ

5. การจำแนกข้อมูลและลักษณะการแจกแจง

การจำแนกข้อมูลและลักษณะการแจกแจง (Classification of Data and Distribution Characteristics) เป็นการพิจารณาว่า **ข้อมูลกระจายตัวอย่างไร** เมื่อวางบนแกนจำนวนและเปรียบเทียบกับ การแจกแจงแบบปกติ (Normal Distribution) การวิเคราะห์นี้มีความสำคัญเพราะช่วยให้นักวิจัยสามารถเลือกใช้ **ค่าสถิติกลาง ค่าสถิติการกระจาย และสถิติอนุมาน** ได้เหมาะสมกับสภาพข้อมูลจริง (สุวิมล ตรีภานันท์, 2546)

1) ประเภทของการแจกแจงข้อมูล

(1) การแจกแจงแบบสมมาตร (Symmetrical Distribution)

- ลักษณะ: กราฟมีรูปร่างสมมาตรทั้งซ้ายและขวา
- คุณสมบัติ: ค่าเฉลี่ย (Mean) $\approx$ มัธยฐาน (Median) $\approx$ ฐานนิยม (Mode)
- ตัวอย่าง: ส่วนสูงของนักเรียนในชั้นเรียนใหญ่ ๆ ที่ไม่มีความแตกต่างมากนัก
- การใช้สถิติ: เหมาะสำหรับสถิติพารามेटริก เช่น t-test, ANOVA, Regression

ตัวอย่าง

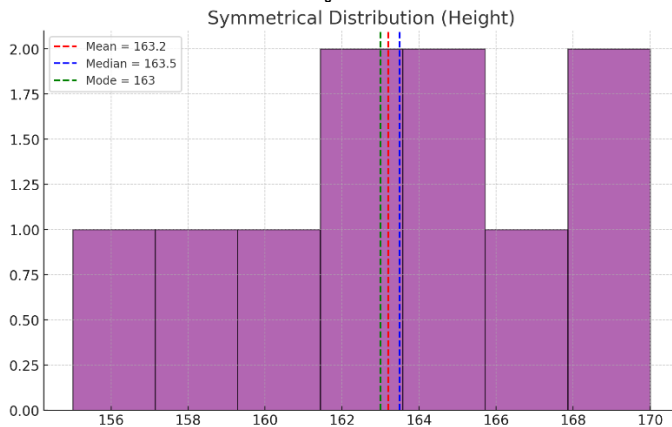
ข้อมูลส่วนสูง (ซม.): 155, 158, 160, 162, 163, 164, 165, 167, 168, 170

- ค่าเฉลี่ย = 163.2

- มัธยฐาน = 163.5

- ฐานนิยม = 163

→ ใกล้เคียงกัน แสดงว่าข้อมูลสมมาตร



แผนภาพที่ 3.10 Symmetrical Distribution

(2) การแจกแจงแบบไม่สมมาตร (Asymmetrical Distribution/ Skewed Distribution)

(ก) การแจกแจงเบ้ขวา (Positive Skew)

- ลักษณะ: ข้อมูลส่วนใหญ่อยู่ค่าต่ำ หางยาวไปทางขวา
- ตัวอย่าง: รายได้ประชากรส่วนใหญ่ต่ำ แต่มีบางคนมีรายได้สูงมาก
- คุณสมบัติ: ค่าเฉลี่ย (Mean) > มัธยฐาน (Median) > ฐานนิยม (Mode)
- การใช้สถิติ: ไม่ควรใช้ Mean เพียงอย่างเดียว → ใช้ Median ร่วม

ตัวอย่าง

รายได้ (พันบาท): 10, 12, 13, 14, 15, 18, 20, 25, 50, 120

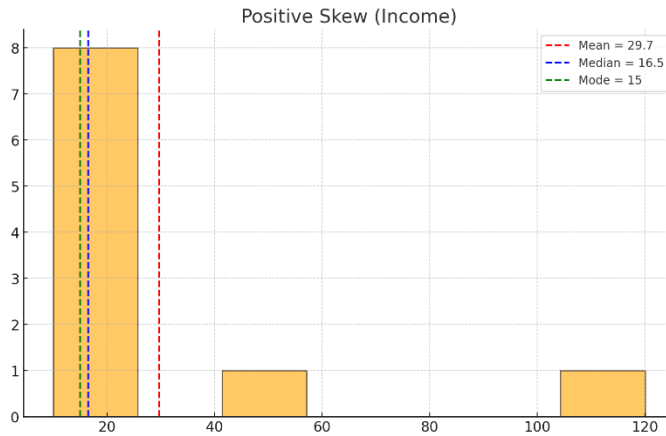
- ค่าเฉลี่ย = 29.7

- มัธยฐาน = 16.5

- ฐานนิยม = 15

→ ค่าเฉลี่ยสูงกว่ามากเพราะถูก “outlier = 120” ดึงขึ้น → การแจกแจง

เบ้ขวา



แผนภาพที่ 3.11 Positive Skew

(ข) การแจกแจงเบ้ซ้าย (Negative Skew)

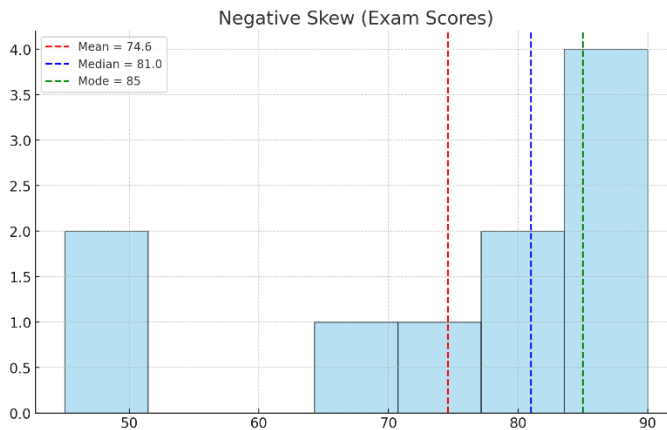
- ลักษณะ: ข้อมูลส่วนใหญ่อยู่ค่ามาก หางยาวไปทางซ้าย
- ตัวอย่าง: คะแนนสอบในห้องเรียนที่ผู้เรียนส่วนใหญ่ทำคะแนนได้ดี แต่มีบางคนสอบตก
- คุณสมบัติ: ค่าเฉลี่ย (Mean) < มัธยฐาน (Median) < ฐานนิยม (Mode)
- การใช้สถิติ: ใช้ Median หรือ Mode เป็นตัวแทนค่ากลางที่เหมาะสมกว่า

ตัวอย่าง

คะแนนสอบ (เต็ม 100): 45, 50, 65, 75, 80, 82, 85, 86, 88, 90

- ค่าเฉลี่ย = 74.6
- มัธยฐาน = 81
- ฐานนิยม = 85

→ ค่าเฉลี่ยถูก “คะแนนต่ำ 45, 50” ดึงลง → การแจกแจงเบ้ซ้าย



แผนภาพที่ 3.12 Negative Skew

2) การเปรียบเทียบกับแจกแจงปกติ (Normal Distribution)

(1) ลักษณะของการแจกแจงปกติ

การแจกแจงปกติ (Normal Distribution) เป็นรูปแบบที่ถือว่าสมบูรณ์แบบที่สุดในเชิงทฤษฎีทางสถิติ

- ลักษณะกราฟ: รูประฆังคว่ำ (Bell-shaped curve) และสมมาตร (Symmetric) รอบค่าเฉลี่ย

- ค่ากลาง: Mean = Median = Mode

- การวัดการเบ้ (Skewness): ≈ 0 (ไม่มีการเอียง)

- การวัดความโด่ง (Kurtosis): ≈ 0 (Mesokurtic – ความโด่งปกติ)

- การกระจาย: ร้อยละ 68 ของข้อมูลอยู่ในช่วง ± 1 SD, ร้อยละ 95 ในช่วง ± 2 SD, และร้อยละ 99.7 ในช่วง ± 3 SD จากค่าเฉลี่ย

ตัวอย่าง: ส่วนสูงหรือผล IQ ของประชากรจำนวนมากมักใกล้เคียงการแจกแจงปกติ

(2) เมื่อการแจกแจง “เบ้” (Skewness $\neq 0$)

- เบ้ขวา (Positive Skew): ข้อมูลส่วนใหญ่อยู่ฝั่งค่าต่ำ แต่มีค่าบางตัวสูงมากดึงหางไปทางขวา

→ ทำให้ค่า Mean > Median > Mode

→ ค่าเฉลี่ยอาจสูงกว่าค่ากลางจริงเพราะถูก outliers ดึงขึ้น

- เบ้ซ้าย (Negative Skew): ข้อมูลส่วนใหญ่อยู่ฝั่งค่าสูง แต่มีค่าบางตัวต่ำมากดึงหางไปทางซ้าย

→ ทำให้ค่า Mean < Median < Mode

→ ค่าเฉลี่ยอาจต่ำกว่าค่ากลางจริงเพราะถูก outliers ดึงลง

ตัวอย่าง:

- เบ้ขวา → รายได้ของประชากร (ส่วนใหญ่รายได้ต่ำ แต่บางคนมีรายได้สูงมาก)
- เบ้ซ้าย → คะแนนสอบที่ผู้เรียนส่วนใหญ่ทำคะแนนได้สูง แต่มีบางคนสอบตก

(3) เมื่อการแจกแจง “โด่ง/แบน” (Kurtosis $\neq 0$)

- Leptokurtic (Kurtosis > 0): กราฟโด่งกว่าปกติ ข้อมูลส่วนใหญ่กระจุกใกล้ค่าเฉลี่ย และหางหนา → มีโอกาสเจอ outliers สูง
- Platykurtic (Kurtosis < 0): กราฟแบนกว่าปกติ ข้อมูลกระจายกว้างหางบาง → โอกาสเจอ outliers ต่ำ
- Mesokurtic (Kurtosis ≈ 0): การแจกแจงใกล้เคียงปกติ

ตัวอย่าง:

- Leptokurtic → คะแนนสอบที่นักเรียนส่วนใหญ่ทำได้ใกล้เคียงกัน แต่มีบางคนได้คะแนนสูงหรือต่ำสุดขีด
- Platykurtic → การกระจายของคะแนนแบบสุ่มกว้าง ๆ ไม่มีการกระจุกกลาง

(4) ความสำคัญในการวิเคราะห์สถิติ**1. ผลต่อการเลือกใช้สถิติอนุมาน**

- การทดสอบทางสถิติแบบพารามेटริก (เช่น t-test, ANOVA, Regression) ตั้งอยู่บนสมมติฐานว่าข้อมูลเป็น Normal Distribution
- ถ้าข้อมูลเบ้หรือมี Kurtosis สูง → ควรใช้ สถิติไม่อิงพารามेटริก (Non-parametric tests) เช่น Mann-Whitney U, Kruskal-Wallis

2. ผลต่อการตีความข้อมูล

- ถ้า Skewness สูง → Mean ไม่ใช่ตัวแทนที่ดีของค่ากลาง ควรใช้ Median
- ถ้า Kurtosis สูง → ต้องระวังค่าผิดปกติที่อาจบิดเบือนผล

3. ผลต่อการวิจัยในสาขาต่าง ๆ

- เศรษฐศาสตร์: การวิเคราะห์รายได้และความเหลื่อมล้ำ มักเจอการแจกแจงเบ้ขวา
- การศึกษา: คะแนนสอบบางครั้งใกล้สมมาตร แต่ในบางกรณีอาจเบ้ซ้าย

- การแพทย์: การกระจายของอายุผู้ป่วยบางโรคอาจบอกว่า เพราะผู้สูงอายุมีมาก

สรุปได้ว่า การจำแนกข้อมูลและการพิจารณาลักษณะการแจกแจงเป็นขั้นตอนสำคัญที่จะช่วยให้การวิเคราะห์ข้อมูลถูกต้องและน่าเชื่อถือ นักวิจัยสามารถเลือกใช้สถิติที่เหมาะสม ป้องกันการตีความผิดพลาด และสื่อสารผลการวิจัยได้อย่างแม่นยำ โดยเฉพาะเมื่อต้องนำไปใช้กับการตัดสินใจเชิงนโยบายหรืองานวิชาการ

6. การนำเสนอข้อมูลด้วยกราฟ

การนำเสนอข้อมูลด้วยกราฟ หมายถึง การใช้สัญลักษณ์เชิงภาพ เช่น แท่งสี่เหลี่ยม วงกลม หรือเส้นโค้ง มาแทนข้อมูลเชิงปริมาณหรือเชิงคุณภาพ เพื่อให้เห็นโครงสร้าง แนวโน้ม และการเปรียบเทียบของข้อมูลได้อย่างชัดเจนและเข้าใจง่ายยิ่งกว่าการดูข้อมูลในรูปตัวเลขเพียงอย่างเดียว (สุวิมล ตรีภานันท์, 2546) กราฟถือเป็นเครื่องมือสื่อสารข้อมูลที่จะช่วยให้การวิเคราะห์สถิติมีความสมบูรณ์ เพราะทำให้ผู้วิจัยและผู้อ่านมองเห็นรูปแบบและความสัมพันธ์ในข้อมูลได้อย่างรวดเร็ว

ความสำคัญการนำเสนอข้อมูลด้วยกราฟ

1. **ทำให้ข้อมูลซับซ้อนเข้าใจง่าย** กราฟช่วยย่อข้อมูลจำนวนมากให้อยู่ในรูปแบบที่ตีความได้ทันที เช่น การแสดงคะแนนสอบนักเรียน 1,000 คน หากใช้เพียงตารางตัวเลขจะยากต่อการทำความเข้าใจ แต่เมื่อใช้อิสโตแกรมหรือกราฟเส้นจะมองเห็นแนวโน้มการกระจายคะแนนได้ชัดเจน

2. **ช่วยให้ผู้อ่านทั่วไปเข้าใจผลการวิเคราะห์ได้ง่ายขึ้น** แม้ผู้อ่านจะไม่มีพื้นฐานทางสถิติ ก็สามารถตีความความแตกต่าง สัดส่วน หรือแนวโน้มได้จากภาพที่สื่อสารชัดเจน

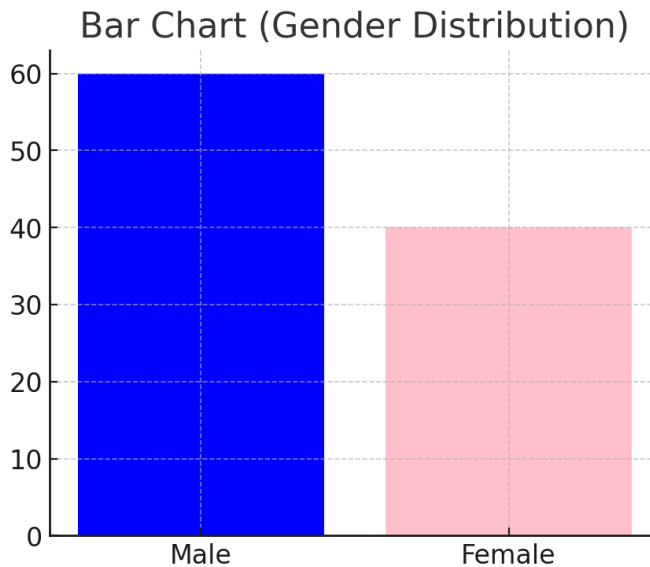
3. **เน้นประเด็นสำคัญที่ผู้วิจัยต้องการสื่อ** เช่น ใช้กราฟแท่งเปรียบเทียบจำนวนผู้ตอบแบบสอบถามในแต่ละจังหวัด หรือใช้กราฟวงกลมแสดงสัดส่วนการใช้เวลาของนักเรียน กราฟเหล่านี้ช่วยให้เห็นสาระสำคัญทันทีโดยไม่ต้องอ่านตัวเลขทั้งหมด

1) ประเภทของกราฟที่ใช้บ่อย

(1) กราฟแท่ง (Bar Chart)

กราฟแท่งเป็นกราฟที่ใช้แท่งสี่เหลี่ยม (Vertical หรือ Horizontal) แสดงค่าความถี่หรือปริมาณ ขนาดความสูงหรือความยาวของแท่งแสดงขนาดของข้อมูลในแต่ละหมวดหมู่

การใช้งาน เหมาะกับข้อมูล **เชิงคุณภาพ (Qualitative Data)** ที่แบ่งเป็นกลุ่ม เช่น เพศ อาชีพ จังหวัด หรือข้อมูลเชิงปริมาณที่จัดเป็นกลุ่มหมวดแล้ว เช่น รายได้แบ่งตามช่วง



แผนภาพที่ 3.13 Bar Chart (Gender Distribution)

ข้อดี

- เข้าใจง่าย เปรียบเทียบค่าระหว่างกลุ่มได้ชัดเจน
- ใช้ได้ทั้งจำนวนจริงและร้อยละ

ข้อจำกัด

- ไม่เหมาะกับข้อมูลต่อเนื่อง
- หากมีหลายหมวดมากเกินไป กราฟอาจรกและอ่านยาก

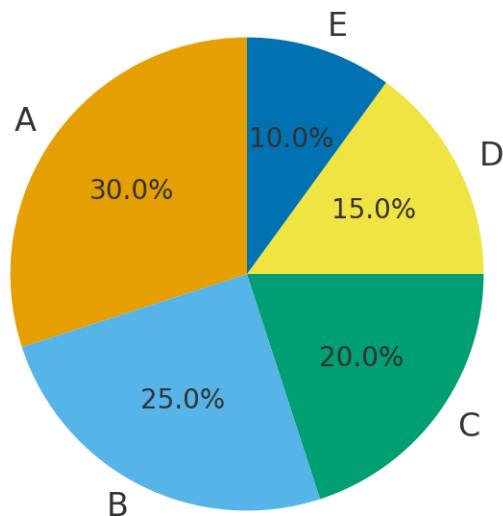
กราฟแท่งเป็นหนึ่งในวิธีที่มีประสิทธิภาพที่สุดในการแสดงการเปรียบเทียบระหว่างหมวดหมู่

(2) กราฟวงกลม (Pie Chart)

แสดงข้อมูลในรูปของวงกลมที่แบ่งออกเป็นส่วน (Sector) แต่ละส่วนมีขนาดตามสัดส่วนหรือร้อยละของข้อมูลทั้งหมด

การใช้งานเหมาะสำหรับการแสดง **สัดส่วนหรือโครงสร้าง** ของข้อมูลที่รวมกันเป็น 100% เช่น การใช้เวลาในหนึ่งวัน หรือสัดส่วนงบประมาณ

Pie Chart (Proportions)



แผนภาพที่ 3.14 Pie Chart (Proportions)

ข้อดี

- ทำให้ผู้อ่านเข้าใจโครงสร้างองค์ประกอบของข้อมูลได้รวดเร็ว
- สื่อถึงแนวคิด “ทั้งหมด = 100%” ได้ชัดเจน

ข้อจำกัด

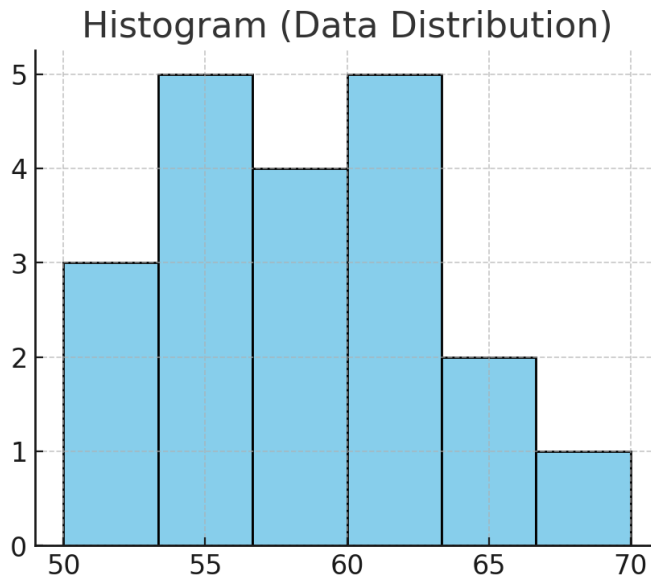
- ถ้ามีชิ้นส่วนมากเกินไป (>6 ส่วน) จะอ่านยาก
- ไม่เหมาะกับการเปรียบเทียบหลายชุดข้อมูล

กราฟวงกลมควรใช้เฉพาะเมื่อข้อมูลมีจำนวนชิ้นส่วนน้อย และต้องการสื่อความหมายเรื่องสัดส่วน

(3) ฮิสโตแกรม (Histogram)

กราฟแท่งที่ใช้แทนการแจกแจงความถี่ของข้อมูลเชิงปริมาณ โดยแต่ละแท่งแทน “ช่วงชั้น” ของข้อมูล

การใช้งาน ใช้สำหรับวิเคราะห์การแจกแจงของข้อมูล เช่น การตรวจสอบว่าข้อมูลมีการกระจายแบบปกติ (Normal Distribution) หรือมีการเบ้ซ้าย/เบ้ขวา



แผนภาพที่ 3.15 Histogram (Data Distribution)

ข้อดี

- ช่วยให้เข้าใจการกระจายของข้อมูลได้ง่าย
- ใช้วิเคราะห์ค่ากลาง ความเบ้ และความโด่งของการแจกแจงได้

ข้อจำกัด

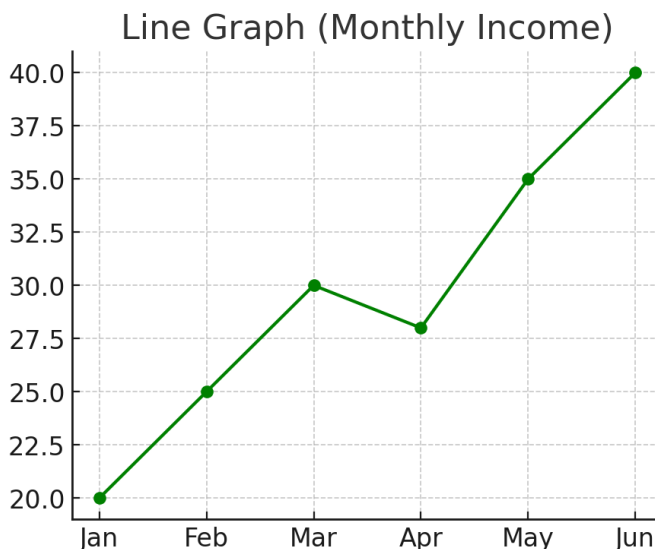
- ผลลัพธ์ขึ้นอยู่กับจำนวนและขนาดของ “ช่วงชั้น” ที่เลือก
- ถ้ากำหนดช่วงชั้นไม่เหมาะสม อาจทำให้ตีความผิด

ฮิสโตแกรมเป็นเครื่องมือหลักในการสำรวจข้อมูลเชิงปริมาณเพื่อหา
ลักษณะการแจกแจง

(4) กราฟเส้น (Line Graph)

ใช้เส้นเชื่อมจุดข้อมูลเรียงตามลำดับ เช่น ตามเวลา เพื่อแสดงการ
เปลี่ยนแปลง

การใช้งาน เหมาะสำหรับการแสดง **แนวโน้ม (Trends)** หรือการ
เปลี่ยนแปลงต่อเนื่อง เช่น รายได้รายเดือน อัตราการว่างงานรายไตรมาส



แผนภาพที่ 3.16 Line Graph (Monthly Income)

ข้อดี

- แสดงการเปลี่ยนแปลงชัดเจน
- เหมาะกับการสังเกตแนวโน้มระยะสั้นและระยะยาว

ข้อจำกัด

- ไม่เหมาะกับข้อมูลที่ไม่มีลำดับเวลา
- ถ้ามีหลายเส้นมากเกินไป กราฟจะซับซ้อนและอ่านยาก

กราฟเส้นเป็นวิธีที่มีประสิทธิภาพสูงสุดในการสื่อสารแนวโน้มและการเปลี่ยนแปลงตามเวลา

2) เกณฑ์การเลือกใช้กราฟที่เหมาะสม

การเลือกใช้กราฟที่เหมาะสมขึ้นอยู่กับ **ชนิดของข้อมูล** และ **วัตถุประสงค์ในการสื่อสาร** หากเลือกกราฟไม่ถูกต้อง อาจทำให้ผู้อ่านตีความผิดหรือข้อมูลขาดความน่าเชื่อถือได้ (Ware, 2013)

(1) เกณฑ์ทั่วไปในการเลือกใช้กราฟ

1. ชนิดข้อมูล

- ข้อมูลเชิงคุณภาพ (Qualitative/Categorical) → เหมาะกับ *Bar Chart* และ *Pie Chart*
- ข้อมูลเชิงปริมาณ (Quantitative/Continuous) → เหมาะกับ *Histogram* และ *Line Graph*

2. เป้าหมายของการนำเสนอ

- เปรียบเทียบ (Comparison): ใช้ *Bar Chart*
- สัดส่วน (Proportion/Composition): ใช้ *Pie Chart* หรือ *Stacked Bar Chart*
- รูปแบบการแจกแจง (Distribution): ใช้ *Histogram*
- แนวโน้มตามเวลา (Trend over time): ใช้ *Line Graph*

3. จำนวนกลุ่มหรือช่วง

- กราฟวงกลมควรใช้เมื่อมี $\leq 5-6$ ส่วน
- กราฟแท่งควรจัดเรียงไม่เกิน $\sim 10-12$ หมวดเพื่อไม่ให้รก
- ฮิสโตแกรมต้องเลือกจำนวนช่วง (bins) ที่เหมาะสมเพื่อให้เห็นรูปทรงจริง

(2) ตัวอย่างการเลือกใช้กราฟ

ตัวอย่างที่ 1 การเปรียบเทียบจำนวนผู้เรียนตามคณะ

ข้อมูล คณะรัฐศาสตร์ 120 คน คณะสังคมศาสตร์ 90 คน

คณะศึกษาศาสตร์ 70 คน

กราฟที่เหมาะสม Bar Chart

เหตุผล ต้องการเปรียบเทียบ “ปริมาณ” ระหว่างคณะ → Bar Chart
สื่อสารได้ชัดเจน

ตัวอย่างที่ 2 การแสดงโครงสร้างการใช้เวลาของนักเรียนต่อวัน

ข้อมูล นอน 8 ชั่วโมง, เรียน 6 ชั่วโมง, ทำการบ้าน 4 ชั่วโมง, พักผ่อน 6 ชั่วโมง

กราฟที่เหมาะสม Pie Chart

เหตุผล เน้น “องค์ประกอบของเวลา 24 ชั่วโมง” ว่าสัดส่วนแต่ละกิจกรรมเป็นเท่าใด

ตัวอย่างที่ 3 วิเคราะห์คะแนนสอบนักเรียน 200 คน

ข้อมูล คะแนนเต็ม 100 กระจายตั้งแต่ 30-95

กราฟที่เหมาะสม Histogram

เหตุผล ข้อมูลต่อเนื่อง ต้องการดูการกระจายว่ามีลักษณะสมมาตรหรือเบ้ → Histogram ช่วยเห็นรูปทรงชัดเจน

ตัวอย่างที่ 4 รายได้เฉลี่ยของครัวเรือนรายเดือน (ม.ค.-ธ.ค.)

ข้อมูล รายได้เพิ่มขึ้นจาก 18,000 บาทในเดือน ม.ค. เป็น 22,000 บาทในเดือน ธ.ค.

กราฟที่เหมาะสม Line Graph

เหตุผล ข้อมูลเป็น time series ต้องการดูแนวโน้ม → Line Graph
แสดงการเปลี่ยนแปลงตามเวลาได้ดีที่สุด

ตารางที่ 3.5 สรุปแนวทางเลือกใช้กราฟ

วัตถุประสงค์	ข้อมูล	กราฟที่แนะนำ
เปรียบเทียบ	เชิงคุณภาพ	กราฟแท่ง (Bar Chart)
แสดงสัดส่วน	เชิงคุณภาพ ($\leq 5-6$ ส่วน)	กราฟวงกลม (Pie Chart)
วิเคราะห์การแจกแจง	เชิงปริมาณต่อเนื่อง	ฮิสโตแกรม (Histogram)
แสดงแนวโน้มตามเวลา	เชิงปริมาณอนุกรมเวลา	กราฟเส้น (Line Graph)

(3) ข้อควรระวัง

- การเลือกใช้กราฟไม่ตรงกับชนิดข้อมูล อาจทำให้ผู้อ่านตีความผิด
- ควรระวังการใช้กราฟวงกลมในกรณีที่มีหลายชิ้นส่วน เพราะการแยกแยะด้วยสายตามีข้อจำกัด (Cleveland & McGill, 1984)

- ต้องใส่ชื่อแกน หน่วย และคำอธิบายประกอบกราฟเสมอ เพื่อป้องกันความคลาดเคลื่อน

3) การตีความข้อมูลจากกราฟ

การตีความข้อมูลจากกราฟ หมายถึง การอ่านข้อความเชิงภาพให้ออกมาเป็นข้อสรุปเชิงสถิติหรือเชิงเนื้อหา นักวิจัยจำเป็นต้องรู้จักวิธีอ่านกราฟแต่ละประเภท เพราะแต่ละชนิดสื่อสารข้อมูลในมิติที่แตกต่างกัน (Few, 2009; Ware, 2013)

(1) แนวทางทั่วไปในการตีความกราฟ

- อ่านหน่วยและสเกลแกนก่อนเสมอ ต้องรู้ว่าแกน X และ Y แทนค่าอะไร หน่วยคืออะไร
- สังเกตความสูง/ขนาดขององค์ประกอบ เช่น ความสูงของแท่ง มุมของวงกลม ความหนาแน่นของฮิสโตแกรม
- หาค่าเด่น จุดที่ข้อมูลกระจุกตัว เช่น ยอดสูงสุดของ Histogram หรือแท่งที่สูงที่สุด
- สังเกตการกระจาย ดูว่าข้อมูลกว้างหรือแคบ หางยาวไปด้านใด
- เปรียบเทียบระหว่างกลุ่ม/ช่วงเวลา หาความต่าง ความคล้าย หรือแนวโน้ม

(2) การตีความกราฟตามประเภท

ก. กราฟแท่ง (Bar Chart)

- **สิ่งที่ดู** ความสูงของแท่งบอกปริมาณ/ความถี่ เปรียบเทียบกันได้ตรง ๆ
- **ตัวอย่าง** กราฟแท่งจำนวนผู้เรียนแบ่งตามช่วงคะแนน
 - แท่งช่วง 60–69 คะแนนสูงที่สุด → นักเรียนส่วนใหญ่อยู่ในช่วงนี้
 - ช่วง 80–89 คะแนนมีน้อยที่สุด → มีนักเรียนจำนวนน้อยที่ทำคะแนนสูงมาก
- **สรุป** กลุ่มคะแนนกลางเป็นกลุ่มใหญ่ที่สุด แสดงว่าระดับความสามารถโดยรวมอยู่ระดับปานกลาง

ข. กราฟวงกลม (Pie Chart)

- **สิ่งที่ดู** สัดส่วนของแต่ละชิ้นส่วน, เปรียบเทียบร้อยละของกลุ่ม
- **ตัวอย่าง** การใช้เวลา 24 ชั่วโมง
 - ทำงาน 41.7% > นอน 33.3% > พักผ่อน 25%
 - ชิ้นส่วนทำงานใหญ่ที่สุด แสดงว่าผู้ตอบใช้เวลามากที่สุดไปกับการทำงาน
- **สรุป** การตีความ Pie Chart มุ่งที่ “สัดส่วน” ไม่ใช่ค่าตัวเลขแน่นอน

ค. ฮิสโตแกรม (Histogram)

- **สิ่งที่ดู** รูปทรงการแจกแจง (ปกติ/เบ้ซ้าย/เบ้ขวา/หลายยอด) จุดสูงสุด (mode) การกระจายตัว
- **ตัวอย่าง** คะแนนสอบ 30 คน
 - Bin 60–70 มีความถี่มากที่สุด (7 คน) → ค่า mode อยู่ช่วงนี้
 - การแจกแจงเบ้ขวาเล็กน้อย เพราะมีคะแนนสูงใกล้ 100
- **สรุป** ข้อมูลส่วนใหญ่กระจุกกลาง แต่มีบางคนได้คะแนนสูงกว่ากลุ่ม

ง. กราฟเส้น (Line Graph)

- **สิ่งที่ดู** แนวโน้มการเปลี่ยนแปลงตามเวลา/ลำดับ, จุดสูงสุด-ต่ำสุด, ความต่อเนื่อง
- **ตัวอย่าง** รายได้ครัวเรือนรายเดือน
 - กราฟเส้นไต่ขึ้นจาก ม.ค. (18,000 บาท) ไป ธ.ค. (22,000 บาท) → รายได้เพิ่มขึ้นต่อเนื่อง
- **สรุป** แนวโน้มเชิงบวก สะท้อนว่ารายได้ครัวเรือนปรับตัวสูงขึ้นตลอดปี

(3) ข้อควรระวังในการตีความ

- **กราฟแท่ง** ถ้าแกน Y ไม่เริ่มที่ 0 ทำให้ค่าต่างเล็ก ๆ ดูแตกต่างมากเกินไปจริง

- **กราฟวงกลม** มุมที่ใกล้กันทำให้ตัดสินใจผิดได้ ไม่ควรใช้เกิน 5-6 ส่วน (Cleveland & McGill, 1984)

- **ฮิสโตแกรม** จำนวน bin ที่เลือกมีผลต่อการตีความ ต้องเลือกอย่างเหมาะสม

- **กราฟเส้น** ไม่ควรเชื่อมจุดข้อมูลที่ไม่มีความต่อเนื่อง (Missing Data) เพราะอาจทำให้เห็นแนวโน้มที่ไม่จริง

สรุปได้ว่า การตีความกราฟคือการอ่านภาพให้เป็นภาษา นักวิจัยต้องเข้าใจทั้งรูปแบบการแสดงผล และบริบทข้อมูล เพื่อหลีกเลี่ยงความเข้าใจผิด การตีความที่ถูกต้องช่วยให้เห็นค่ากลาง การกระจาย แนวโน้ม และความแตกต่างของข้อมูลได้ชัดเจนยิ่งขึ้น

4) ข้อควรระวังในการใช้กราฟ

แม้ว่ากราฟจะเป็นเครื่องมือที่มีประสิทธิภาพในการนำเสนอข้อมูล แต่หากใช้ไม่ถูกต้องอาจทำให้ข้อมูลบิดเบือนและผู้อ่านตีความผิดได้ (Cairo, 2016; Kosslyn, 2006) ดังนั้น ผู้วิจัยจำเป็นต้องตระหนักถึงข้อควรระวัง ดังนี้

(1) การเลือกสเกลแกน (Axis Scale) ไม่เหมาะสม

- **ปัญหา** หากแกน Y ไม่เริ่มที่ศูนย์ (0) อาจทำให้ความแตกต่างเล็กน้อยระหว่างกลุ่มดูใหญ่เกินจริง

- **ตัวอย่าง** คะแนนเฉลี่ย 70 กับ 72 ถ้าแกน Y เริ่มที่ 65 จะดูต่างกันมาก แต่ถ้าเริ่มที่ 0 ความต่างแทบไม่ชัด

- **ข้อแนะนำ** เริ่มแกน Y ที่ศูนย์เสมอ ยกเว้นกรณีข้อมูลมีความละเอียดสูง และต้องใส่หมายเหตุให้ชัด

(2) การใช้กราฟวงกลมที่มีหลายชิ้นเกินไป

- **ปัญหา** การแสดงสัดส่วนมากกว่า 5-6 ส่วน ทำให้สายตามนุษย์เปรียบเทียบพื้นที่/มุมได้ยาก (Cleveland & McGill, 1984)

- **ตัวอย่าง** Pie Chart ที่มี 10 ส่วนขึ้นไป ผู้อ่านไม่สามารถบอกได้ว่าส่วนใดใหญ่กว่ากัน

- **ข้อแนะนำ** ใช้ *Bar Chart* หรือ *Stacked Bar Chart* แทนเมื่อมีหลายหมวดหมู่

(3) การใช้เอฟเฟกต์ 3D หรือการตกแต่งเกินจำเป็น

- **ปัญหา** กราฟ 3D ทำให้ความสูงหรือพื้นที่ถูกบิดเบือน การตีความผิดพลาดได้ง่าย (Tufte, 2001)

- **ตัวอย่าง** Bar Chart 3D ที่มุมมองทำให้แท่งหน้าดูใหญ่กว่าความจริง

- **ข้อแนะนำ** ใช้กราฟ 2D ธรรมดาเพื่อความชัดเจน และเน้นความถูกต้องมากกว่าความสวยงาม

(4) การใช้สีที่ไม่เหมาะสม

- ปัญหา สีที่ใกล้เคียงกันเกินไปอาจทำให้แยกยาก โดยเฉพาะสำหรับผู้ที่มีปัญหาสีตาบอด (color blindness)
- ตัวอย่าง ใช้สีเขียวอ่อน-เขียวเข้มใน Pie Chart หลายส่วน ยากต่อการจำแนก
- ข้อเสนอแนะ ใช้สีที่มีความแตกต่างชัดเจน และควรมีคำอธิบายหรือ Legend กำกับเสมอ

5 การละเลยคำอธิบายกราฟ

- ปัญหา กราฟที่ไม่มีชื่อ (Title), ชื่อแกน (Axis Label), หน่วย (Unit) หรือคำอธิบาย (Legend) ทำให้ข้อมูลขาดความหมาย
- ข้อเสนอแนะ ใส่ส่วนประกอบสำคัญครบถ้วน ได้แก่
 - ชื่อกราฟ (Title)
 - ชื่อแกนและหน่วย (Axis labels with units)
 - คำอธิบายสัญลักษณ์ (Legend)
 - แหล่งที่มาของข้อมูล (Source)

6 การตีความเกินจริงจากกราฟ

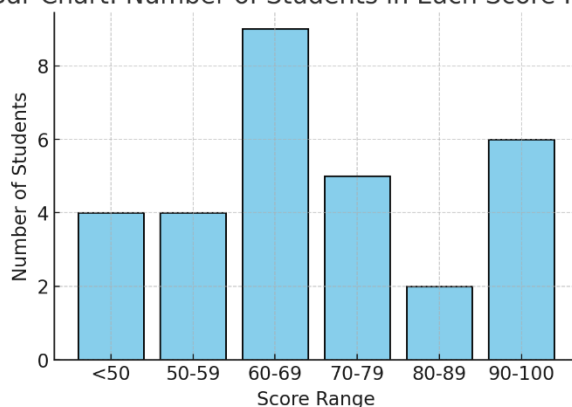
- ปัญหา ผู้นำเสนออาจเลือกช่วงเวลา/กลุ่มเฉพาะเพื่อเน้นผลลัพธ์บางอย่าง (Selective Presentation)
 - ตัวอย่าง แสดงข้อมูลรายได้เฉพาะช่วง 3 เดือนที่เพิ่มขึ้น ทั้งที่ทั้งปีมีการขึ้นลงสลับกัน
 - ข้อเสนอแนะ ต้องนำเสนอข้อมูลครบถ้วน และอธิบายข้อจำกัดอย่างโปร่งใส
- สรุปได้ว่า ข้อควรระวังในการใช้กราฟเน้นที่ ความถูกต้องและความชัดเจนมากกว่าความสวยงาม หากเลือกสเกลผิด ใช้ 3D ตกแต่งมากเกินไป หรือใช้ Pie Chart ที่ซับซ้อน จะทำให้ผู้อ่านเข้าใจผิดได้ ดังนั้น กราฟที่ดีควรมีรูปแบบเรียบง่าย มีการอธิบายองค์ประกอบครบ และสะท้อนข้อมูลอย่างตรงไปตรงมา

5) ตัวอย่างประกอบ (จากคะแนนสอบนักเรียน 30 คน)

ตารางที่ 3.6 ข้อมูลสรุป (ช่วงคะแนนแบบ Bar/Pie)

ช่วงคะแนน	ความถี่ (f)	ร้อยละ (%)
ต่ำกว่า 50	4	13.3
50-59	4	13.3
60-69	9	30.0
70-79	5	16.7
80-89	2	6.7
90-100	6	20.0

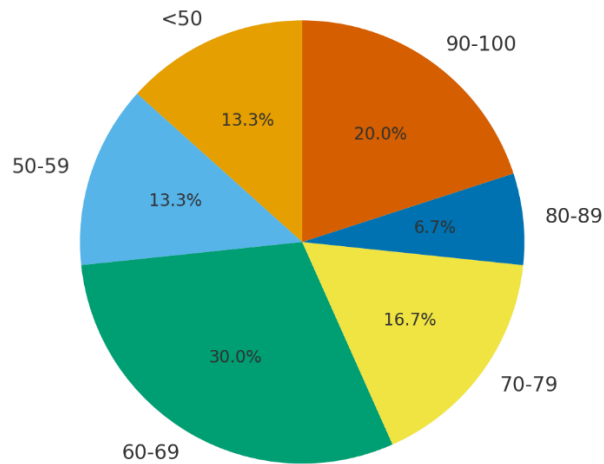
Bar Chart: Number of Students in Each Score Range



แผนภาพที่ 3.17 Bar Chart

แผนภาพที่ 3.17 กราฟแท่งแสดงจำนวนผู้เรียนตามช่วงคะแนนสอบ นักเรียนส่วนใหญ่อยู่ในช่วงคะแนน 60-69 คะแนน (9 คน หรือ 30.0%) ซึ่งเป็นแท่งที่สูงที่สุด

Pie Chart: Proportion of Students by Score Range



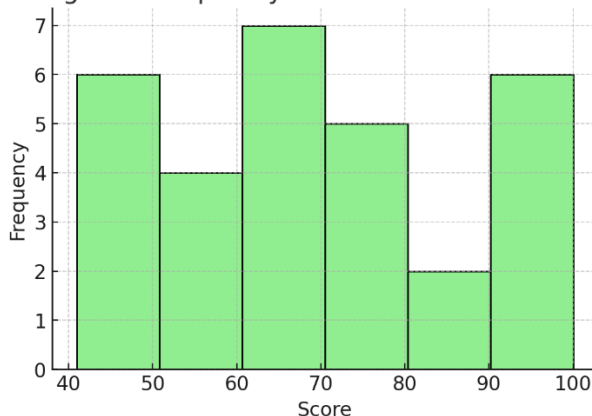
แผนภาพที่ 3.18 Pie Chart

แผนภาพที่ 3.18 กราฟวงกลมแสดงสัดส่วนผู้เรียนตามช่วงคะแนนสอบ พบว่า ช่วงคะแนน 60–69 มีสัดส่วนมากที่สุด (30.0%) รองลงมาคือช่วง 90–100 (20.0%)

ตารางที่ 3.7 Histogram (bins = 6 จาก min–max)

Bin	ความถี่ (f)	ความกว้างชั้น (i)	ความหนาแน่น (d)
[41.0, 50.8)	6	9.83	0.0204
[50.8, 60.7)	4	9.83	0.0136
[60.7, 70.5)	7	9.83	0.0237
[70.5, 80.3)	5	9.83	0.0170
[80.3, 90.2)	2	9.83	0.0068
[90.2, 100.0]	6	9.83	0.0204

Histogram: Frequency Distribution of Exam Scores



แผนภาพที่ 3.19 Histogram

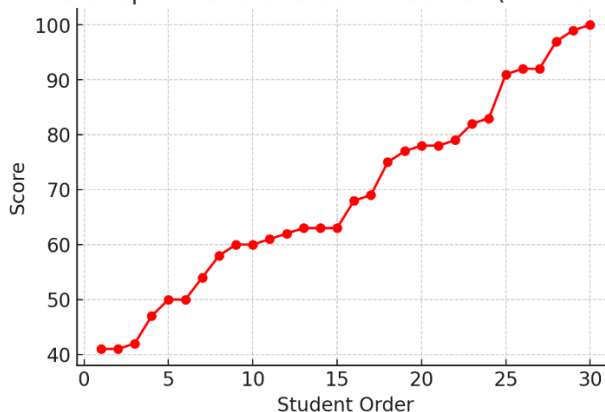
แผนภาพที่ 3.19 ฮิสโตแกรมแสดงการแจกแจงคะแนนสอบของนักเรียน 30 คน โดยใช้ 6 ช่วงชั้น พบว่าคะแนนกระจุกตัวอยู่ในช่วง 60-70 คะแนน และมีการกระจายไปทางคะแนนสูง (เบ้ขวาเล็กน้อย)

ตาราง Line Graph (คะแนนเรียงจากน้อย→มาก)

5 ค่าต่ำสุด: 41, 41, 42, 47, 50

5 ค่าสูงสุด: 92, 92, 97, 99, 100

Line Graph: Trend of Student Scores (Low to High)



แผนภาพที่ 3.20 Line Graph

แผนภาพที่ 3.20 กราฟเส้นแสดงคะแนนสอบนักเรียน 30 คนเรียงจากน้อยไปมาก พบว่ามีการได้ระดับอย่างต่อเนื่องจาก 41 คะแนน ไปจนถึง 100 คะแนน โดยมีช่วงสูงสุดอยู่ที่ 97-100 คะแนน

สรุปได้ว่า การนำเสนอข้อมูลด้วยกราฟเป็นกระบวนการที่สำคัญในการวิเคราะห์เชิงสถิติ เนื่องจากช่วยย่อข้อมูลจำนวนมากให้เข้าใจง่าย สื่อสารได้ชัดเจน และตีความได้รวดเร็ว โดยกราฟแต่ละประเภทมีข้อดี ข้อจำกัด และเหมาะกับวัตถุประสงค์ที่ต่างกัน เช่น กราฟแท่งเหมาะกับการเปรียบเทียบ กราฟวงกลมเหมาะกับการแสดงสัดส่วน ฮิสโตแกรมใช้ตรวจสอบการแจกแจง และกราฟเส้นแสดงแนวโน้มตามเวลา การเลือกใช้กราฟที่เหมาะสมและระมัดระวังข้อผิดพลาดในการนำเสนอ จะทำให้การสื่อสารข้อมูลมีความถูกต้อง โปร่งใส และช่วยให้นักวิจัยกับผู้อ่านเข้าใจลักษณะการกระจาย แนวโน้ม และโครงสร้างของข้อมูลได้อย่างสมบูรณ์

สรุปท้ายบท

การแจกแจงข้อมูลและสถิติเชิงพรรณนา ได้อธิบายแนวคิดและเครื่องมือพื้นฐานที่สำคัญสำหรับการทำความเข้าใจข้อมูลเชิงสถิติอย่างเป็นระบบ โดยเริ่มจากตารางแจกแจงความถี่ ซึ่งเป็นวิธีการจัดระเบียบข้อมูลจำนวนมากให้อยู่ในรูปแบบที่เข้าใจง่ายและช่วยให้เห็นลักษณะการกระจายอย่างชัดเจน ต่อมาคือค่ากลาง ได้แก่ ค่าเฉลี่ย มัธยฐาน และฐานนิยม ซึ่งใช้เป็นตัวแทนแนวโน้มของข้อมูลแต่ละชุด อย่างไรก็ตาม ค่ากลางเพียงอย่างเดียวไม่สามารถอธิบายความแตกต่างของข้อมูลได้ จึงต้องใช้ค่าการกระจาย เช่น พิสัย ส่วนเบี่ยงเบนมาตรฐาน และพิสัยควอไทล์ มาช่วยอธิบายความแปรปรวนหรือความแตกต่างภายในชุดข้อมูล

นอกจากนี้ การวิเคราะห์รูปร่างของการแจกแจงด้วยค่า Skewness และ Kurtosis ยังช่วยให้เข้าใจลักษณะการเอียงและความโด่งหรือความแบนของการกระจาย ซึ่งมีผลต่อการเลือกใช้สถิติอนุมานที่เหมาะสม การจำแนกข้อมูลและลักษณะการแจกแจงยังทำให้นักวิจัยสามารถเปรียบเทียบกับการแจกแจงปกติ และตัดสินใจเลือกใช้ค่ากลางหรือวิธีวิเคราะห์ที่ตรงกับลักษณะข้อมูลได้อย่างถูกต้อง สุดท้าย การนำเสนอข้อมูลด้วยกราฟ ไม่ว่าจะเป็นกราฟแท่ง กราฟวงกลม ฮิสโตแกรม หรือกราฟเส้น ช่วยให้การสื่อสารข้อมูลชัดเจนขึ้น ทำให้ผู้อ่านเข้าใจแนวโน้ม การเปรียบเทียบ และโครงสร้างของข้อมูลได้รวดเร็วและตรงประเด็น

กล่าวโดยสรุป สถิติเชิงพรรณนามีบทบาทสำคัญอย่างยิ่งในกระบวนการวิเคราะห์ข้อมูล เพราะเป็นขั้นตอนพื้นฐานที่ช่วยย่อข้อมูลจำนวนมากให้อยู่ในรูปแบบที่ง่ายต่อการทำความเข้าใจ ช่วยให้ผู้วิจัยเห็นภาพรวมของข้อมูลอย่างถูกต้อง และเป็นรากฐานในการนำไปสู่การวิเคราะห์เชิงลึกและการใช้สถิติขั้นสูงต่อไปได้อย่างมีประสิทธิภาพ

บทที่ 4

ความน่าจะเป็นและการแจกแจงความน่าจะเป็น (PROBABILITY AND PROBABILITY DISTRIBUTIONS)

ความน่าจะเป็นและการแจกแจงความน่าจะเป็นเป็นพื้นฐานสำคัญของการศึกษาทางสถิติ โดยเฉพาะในเชิงอนุมาน เนื่องจากความน่าจะเป็นทำหน้าที่เป็น “เครื่องมือทางคณิตศาสตร์” ที่ช่วยให้นักวิจัยสามารถอธิบายและตัดสินใจได้อย่างมีเหตุผลภายใต้ความไม่แน่นอน (Casella & Berger, 1990; Walpole, 1970) การทำความเข้าใจแนวคิดของความน่าจะเป็นจึงเป็นจุดเริ่มต้นที่สำคัญในการต่อยอดไปสู่การสร้างแบบจำลองทางสถิติ การทดสอบสมมติฐาน การสร้างช่วงความเชื่อมั่น และการคาดการณ์พฤติกรรมของประชากรในสังคม

ในงานวิจัยทางสังคมศาสตร์ ความน่าจะเป็นมีบทบาทอย่างมากในการอธิบายข้อมูลจากการสุ่มตัวอย่าง ซึ่งไม่สามารถเก็บจากประชากรทั้งหมดได้ ตัวอย่างเช่น การสำรวจความคิดเห็นทางการเมือง หรือการวิเคราะห์พฤติกรรมผู้บริโภค ล้วนต้องอาศัยแนวคิดความน่าจะเป็นเพื่อประมาณโอกาสและระดับความน่าเชื่อถือของผลลัพธ์ที่ได้ (Milton & Arnold, 1995)

นอกจากนี้ การแจกแจงความน่าจะเป็น (Probability Distributions) ยังช่วยให้เข้าใจรูปแบบการกระจายของข้อมูล ซึ่งแบ่งได้เป็น การแจกแจงแบบไม่ต่อเนื่อง เช่น การแจกแจงทวินาม (Binomial) และการแจกแจงปัวส์ซอง (Poisson) ที่ใช้กับข้อมูลการนับ และการแจกแจงแบบต่อเนื่อง เช่น การแจกแจงปกติ (Normal) ที่ใช้กับข้อมูลต่อเนื่องอย่างคะแนนสอบหรือรายได้ ซึ่งถือเป็นเครื่องมือสำคัญที่นักวิจัยใช้ในการตีความและวิเคราะห์ปรากฏการณ์ทางสังคม (Dudewicz, 1988)

กล่าวโดยสรุป บทนี้จะอธิบายความหมายและแนวคิดของความน่าจะเป็น กฎของความน่าจะเป็น การแจกแจงแบบไม่ต่อเนื่องและต่อเนื่อง การใช้ตาราง Z และความสำคัญของความน่าจะเป็นต่อการวิเคราะห์เชิงอนุมาน เพื่อวางรากฐานในการทำความเข้าใจสถิติอนุมานและการประยุกต์ใช้ในงานวิจัยทางสังคมศาสตร์อย่างมีประสิทธิภาพ

1. ความหมายและแนวคิดของความน่าจะเป็น

ความน่าจะเป็น (Probability) เป็นหัวใจสำคัญของวิชาสถิติ ใช้สำหรับบอก “โอกาส” หรือ “ความเป็นไปได้” ที่เหตุการณ์ใดเหตุการณ์หนึ่งจะเกิดขึ้น โดยมีค่าอยู่ในช่วง 0 ถึง 1 ซึ่ง 0 หมายถึงไม่มีทางเกิดขึ้น และ 1 หมายถึงต้องเกิดขึ้นแน่นอน (Casella & Berger, 1990; ทวีรัตน์ ศิวคุลย์, 2539) ตัวอย่างเช่น ความน่าจะเป็นของการออกหัวเมื่อโยนเหรียญที่สมดุลคือ 0.5 หรือ 50%

1) ความหมายและสัญลักษณ์

ถ้าให้ A แทนเหตุการณ์ที่สนใจ ความน่าจะเป็นที่จะเกิดเหตุการณ์ A เขียนแทนด้วย $P(A)$ และหาค่า $P(A)$ ได้จาก

$$P(A) = \frac{\text{จำนวนผลลัพธ์ของเหตุการณ์}}{\text{จำนวนผลลัพธ์ที่เป็นไปได้ทั้งหมด}}$$

วิธีการหาจำนวนผลลัพธ์ของเหตุการณ์ที่สนใจ และจำนวนผลลัพธ์ที่เป็นไปได้ทั้งหมดมีสองวิธี ซึ่งเราจะศึกษาการหาจำนวนผลลัพธ์ทั้งสองจากตัวอย่างเช่น ถ้า A หมายถึงเหตุการณ์ที่ “นักศึกษามาถึงห้องเรียนก่อนเวลา 8.30 น.” และจากข้อมูลพบว่า มีนักศึกษา 60 จาก 72 คนที่มาถึงก่อนเวลา เราสามารถคำนวณได้ว่า

$$P(A) = \frac{60}{72} = 0.833$$

หมายถึงความน่าจะเป็นที่นักศึกษามาถึงก่อนเวลาเท่ากับ 83.3%

การตีความค่า $P(A)$

- ถ้า $P(A) = 1 \rightarrow$ เกิดขึ้นแน่นอน
- ถ้า $P(A) = 0 \rightarrow$ ไม่มีทางเกิดขึ้น
- ถ้า $P(A) = 0.7 \rightarrow$ มีโอกาสเกิดขึ้น 70%

2) วิธีการหาความน่าจะเป็น

วิธีการหาความน่าจะเป็นมีทั้งแบบเชิงทฤษฎีและเชิงปฏิบัติ

2.1) วิธีการแบบคลาสสิก (Classical Method)

อิงจากหลักการนับจำนวนผลลัพธ์ที่เท่ากันทุกกรณี เช่น การโยนเหรียญหนึ่งครั้งมีผลลัพธ์สองแบบคือหัวหรือก้อย ดังนั้น $P(\text{ออกหัว}) = 1/2 = 0.5$ หรือ 50%

2.2) วิธีการเชิงประจักษ์ (Empirical Method)

ใช้ข้อมูลจากการทดลองซ้ำหรือสถิติในอดีต เช่น การโยนเหรียญ 100 ครั้งพบว่า ออกหัว 80 ครั้ง จะได้ $P(\text{ออกหัว}) = 0.80$ ซึ่งสะท้อนสภาพจริงมากกว่าทฤษฎี (Milton & Arnold, 1995)

3) คุณสมบัติของความน่าจะเป็น

ให้ A และ B เป็นเหตุการณ์ใด ๆ คุณสมบัติของความน่าจะเป็น คือ

1. $0 \leq P(A) \leq 1$
2. $P(\text{ผลลัพธ์ทั้งหมด}) = 1$ และ $P(\text{ไม่เกิดผลลัพธ์เลย}) = 0$
3. $P(\text{ไม่เกิด } A) = 1 - P(A)$
4. ถ้าเหตุการณ์ A และ B ไม่เกิดร่วมกัน $\rightarrow P(A \cup B) = P(A) + P(B)$
5. ถ้าเหตุการณ์ A และ B เกิดร่วมกันได้ $\rightarrow P(A \cup B) = P(A) + P(B) -$

$P(A \cap B)$ (ทวิรัตน์ ศิริกุลย์, 2539)

4) ตัวแปรสุ่มและการแจกแจงความน่าจะเป็น

ผลลัพธ์จากการทดลองสุ่มมักแทนด้วย ตัวแปรสุ่ม (Random Variable) ซึ่งแบ่งเป็น 2 ประเภท (สมจิต วัฒนาชยากุล, 2546)

(1) **ตัวแปรสุ่มไม่ต่อเนื่อง (Discrete Random Variable)** มีค่าที่นับได้ เช่น จำนวนครั้งที่นักศึกษาเข้าร่วมกิจกรรมชมรม

(2) **ตัวแปรสุ่มต่อเนื่อง (Continuous Random Variable)** มีค่าที่เป็นช่วงของจำนวนจริง เช่น รายได้ ความสูง หรือคะแนนสอบ

การแจกแจงความน่าจะเป็น (Probability Distribution) คือฟังก์ชันที่บอกความน่าจะเป็นของแต่ละค่าที่ตัวแปรสุ่มสามารถรับได้

5) ความสำคัญของความน่าจะเป็นในสังคมศาสตร์

ในงานวิจัยทางสังคมศาสตร์ ความน่าจะเป็นช่วยให้นักวิจัยสามารถ:

- ประเมิน **ความเสี่ยง** ของนโยบายหรือโครงการ
- ทำนาย **พฤติกรรม** ของประชาชนจากการสุ่มตัวอย่าง
- ใช้ในการ **ทดสอบสมมติฐาน** และ **สร้างช่วงความเชื่อมั่น**
- พัฒนา **แบบจำลองการแจกแจง** เพื่อตีความข้อมูลจากการสำรวจ (มนตรี

สังข์ทอง, 2557; Milton & Arnold, 1995)

ตัวอย่างเช่น การคาดการณ์ผลเลือกตั้งโดยใช้ข้อมูลจากการสำรวจตัวอย่างผู้มีสิทธิเลือกตั้ง หรือการประเมินทัศนคติประชาชนต่อการจัดการสิ่งแวดล้อม โดยอาศัยแนวคิดความน่าจะเป็นในการประมาณโอกาสและความน่าเชื่อถือของผลที่ได้

สรุปได้ว่า ความน่าจะเป็นไม่ใช่เพียงศาสตร์ทางคณิตศาสตร์ แต่เป็นภาษาที่ใช้สื่อสาร “ความไม่แน่นอน” ในการวิจัยทางสังคมศาสตร์ การเข้าใจความหมาย วิธีการหาค่าคุณสมบัติ ตัวแปรสุ่ม และแบบจำลองการแจกแจง จะช่วยให้ นักวิจัยสามารถตีความวิเคราะห์ และตัดสินใจได้อย่างเป็นระบบและน่าเชื่อถือ

6) แนวคิดหลักของความน่าจะเป็น

ความน่าจะเป็นสามารถตีความได้หลายวิธี ขึ้นอยู่กับมุมมองและการประยุกต์ใช้ โดยแนวคิดหลักที่นิยมใช้กันทั่วไปมี 3 แบบ ดังนี้

(1) ความน่าจะเป็นเชิงทฤษฎี (Theoretical Probability)

ความน่าจะเป็นเชิงทฤษฎีเกิดจากการอธิบายตามหลักคณิตศาสตร์บริสุทธิ์ โดยใช้นับจำนวนกรณีที่เป็นไปได้และกรณีที่สนใจ ตัวอย่างเช่น การทอยลูกเต๋าหนึ่งลูก มี 6 หน้าเท่ากันทุกหน้า ดังนั้นโอกาสที่จะได้เลข 4 เท่ากับ

$$P(\text{ได้เลข 4}) = \frac{1}{6}$$

แนวคิดนี้ถือว่าทุกผลลัพธ์มีโอกาสเท่ากัน เหมาะกับสถานการณ์ที่ควบคุมได้ดี และผลลัพธ์เป็นแบบสมมาตร เช่น การโยนเหรียญ การจับสลาก หรือการเล่นไพ่

ตัวอย่างการประยุกต์ในสังคมศาสตร์: การสุ่มเลือกตัวอย่างประชากรจากกลุ่มที่มีโอกาสเท่ากันทุกคน เช่น การสุ่มตัวอย่างนักศึกษา 1 คนจากรายชื่อ 100 คน ความน่าจะเป็นที่จะเลือกนักศึกษาใดนักศึกษาคือ $1/100$

(2) ความน่าจะเป็นเชิงประจักษ์ (Empirical Probability)

เป็นการหาความน่าจะเป็นจากข้อมูลจริงหรือจากการทดลองซ้ำ ๆ หลายครั้ง โดยใช้สัดส่วนของความถี่สัมพัทธ์เป็นตัวประมาณค่า เช่น การโยนเหรียญ 100 ครั้ง ได้หัว 48 ครั้ง ดังนั้นโอกาสที่เหรียญจะออกหัวคือ

$$P(\text{หัว}) = \frac{48}{100} = 0.48$$

แนวคิดนี้สะท้อนความเป็นจริงมากกว่า เพราะอาจพบความคลาดเคลื่อนจากทฤษฎี เช่น เหรียญอาจไม่สมมาตร หรือการโยนมีแรงไม่เท่ากัน

ตัวอย่างการประยุกต์ในสังคมศาสตร์: การวิเคราะห์สถิติอาชญากรรม ถ้าในพื้นที่หนึ่งพบเหตุลักทรัพย์ 50 ครั้งจากคดีทั้งหมด 200 ครั้งในปีที่ผ่านมา ความน่าจะเป็นเชิงประจักษ์ของการเกิดเหตุลักทรัพย์คือ $50/200 = 0.25$ หรือ 25%

(3) ความน่าจะเป็นเชิงอัตนัย (Subjective Probability)

ความน่าจะเป็นเชิงอัตนัยขึ้นอยู่กับ “การตัดสินใจหรือความเชื่อ” ของบุคคล อาจพิจารณาจากประสบการณ์ ความรู้ หรือข้อมูลส่วนบุคคล ตัวอย่างเช่น แพทย์ประเมินว่าโอกาสที่ผู้ป่วยจะหายจากการรักษาอยู่ที่ 70% แม้จะไม่ได้มีการคำนวณเชิงสถิติที่ชัดเจน แต่สะท้อน “ระดับความมั่นใจ” ของผู้เชี่ยวชาญ

ตัวอย่างการประยุกต์ในสังคมศาสตร์: นักวิเคราะห์การเมืองอาจคาดการณ์ว่า ผู้สมัคร A มีโอกาสชนะการเลือกตั้ง 65% โดยพิจารณาจากแนวโน้มคะแนนเสียง การสนับสนุนจากพรรคการเมือง และประสบการณ์การเลือกตั้งที่ผ่านมา

สรุปได้ว่า

- ทฤษฎี (Theoretical) อิงหลักการคณิตศาสตร์ ใช้เมื่อผลลัพธ์สมมาตรและควบคุมได้

- ประจักษ์ (Empirical) อิงข้อมูลและการทดลองจริง ใช้เมื่อต้องการสะท้อนสถานการณ์ตามข้อเท็จจริง

- อัตนัย (Subjective) อิงประสบการณ์และความเชื่อ ใช้ในกรณีที่ข้อมูลไม่ครบถ้วนหรือขึ้นกับการตัดสินใจของผู้เชี่ยวชาญ

ทั้งสามแนวคิดนี้ล้วนมีบทบาทในงานวิจัยทางสังคมศาสตร์ โดยเฉพาะเมื่อผลสับสน เช่น การคาดการณ์ผลการเลือกตั้งอาจใช้ข้อมูลเชิงประจักษ์ (ผลสำรวจ) ผสมกับทฤษฎีการสุ่มตัวอย่าง และการตีความเชิงอัตนัยของนักวิเคราะห์การเมือง

2. กฎของความน่าจะเป็น

กฎของความน่าจะเป็น เป็นเครื่องมือพื้นฐานสำหรับคำนวณโอกาสของเหตุการณ์ซับซ้อน โดยเฉพาะเมื่อเราต้องรวม (union) หรือซ้อนทับ (intersection) เหตุการณ์หลายเหตุการณ์เข้าด้วยกัน (Casella & Berger, 1990; Walpole, 1970)

ก่อนเริ่ม ใช้สัญลักษณ์มาตรฐาน

- A, B = เหตุการณ์ (events)
- A^c = เหตุการณ์ตรงข้ามของ A (complement)
- $P(\cdot)$ = ความน่าจะเป็น
- $\cup$ = “หรือ” (union), $\cap$ = “และ/พร้อมกัน” (intersection)

หมายเหตุ: กฎสองข้อต่อไปนี้จะทำงานคู่กับความน่าจะเป็นแบบมีเงื่อนไข $P(B | A)$ และแนวคิดความเป็นอิสระของเหตุการณ์ เพื่อหาความน่าจะเป็นของเหตุการณ์ซ้อนทับใช้กฎพื้นฐาน 2 ประการ ดังนี้

1) กฎบวก (Addition Rule)

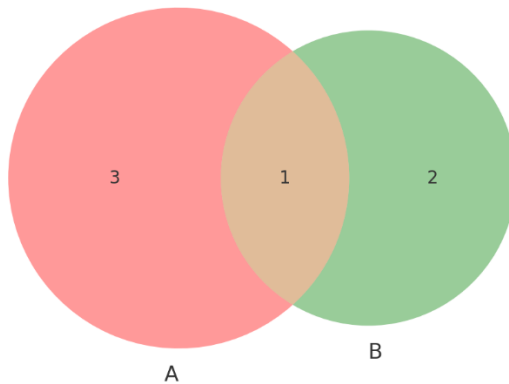
ใช้เมื่อถามความน่าจะเป็นของเหตุการณ์ “อย่างน้อยหนึ่งเหตุการณ์เกิดขึ้น” หรือ “ A หรือ B เกิด”

รูปทั่วไป (มีการทับซ้อน)

$$P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

แนวคิดคือ “รวมพื้นที่ของ A และ B แล้วลบส่วนที่นับซ้ำ” ตรงบริเวณซ้อนทับ (Venn diagram) (Walpole, 1970)

Addition Rule: $P(A \cup B) = P(A) + P(B) - P(A \cap B)$



แผนภาพที่ 4.1 Venn Diagram

กรณีไม่ทับซ้อน (Mutually Exclusive)

ถ้า $A \cap B = \emptyset$ (เกิดพร้อมกันไม่ได้) $\Rightarrow P(A \cup B) = P(A) + P(B)$

(Casella & Berger, 1990)

ตัวอย่าง (ลูกเต๋า 1 ลูก)

ให้ $A = \{\text{ได้เลขคู่}\} = \{2, 4, 6\} \Rightarrow P(A) = 3/6$ และ

$B = \{\text{ได้เลขมากกว่า 4}\} = \{5, 6\} \Rightarrow P(B) = 2/6$

ส่วนที่ทับซ้อน $A \cap B = \{6\} \Rightarrow P(A \cap B) = 1/6$

ดังนั้น

$$P(A \cup B) = \frac{3}{6} + \frac{2}{6} - \frac{1}{6} = \frac{4}{6} = \frac{2}{3}$$

ตัวอย่างทางสังคมศาสตร์ (สำรวจทัศนคติ)

สมมติจากโพลเมืองหนึ่ง

- $P(A) = 0.55$ = โอกาสที่ผู้ตอบเป็น “เพศหญิง”

- $P(B) = 0.62$ = โอกาสที่ “เห็นด้วยนโยบาย X”

- $P(A \cap B) = 0.40$ = โอกาสที่ “เป็นเพศหญิงและเห็นด้วยนโยบาย X”

ต้องการ $P(A \cup B)$ = โอกาสที่ “เป็นเพศหญิงหรือเห็นด้วยนโยบาย X (หรือทั้งสองอย่าง)”

$$P(A \cup B) = 0.55 + 0.62 - 0.40 = 0.77$$

แปลว่า 77% ของประชากรมีอย่างน้อยหนึ่งคุณลักษณะตามที่ระบุ (Milton & Arnold, 1995)

เคล็ดลับที่ใช้บ่อย “กฎเสริม” (Complement Rule)

$$P(A^c) = 1 - P(A)$$

ใช่ง่ายมากเวลาอยากหา “อย่างน้อย 1 ครั้ง” $\Rightarrow$ คิด “ไม่เกิดเลยสักครั้ง” แล้วนำไปลบจาก 1 (Casella & Berger, 1990)

ขยายสำหรับหลายเหตุการณ์ (สามเหตุการณ์ขึ้นไป)

$$\begin{aligned} P(A \cup B \cup C) &= P(A) + P(B) + P(C) \\ &\quad - P(A \cap B) - P(A \cap C) - P(B \cap C) \\ &\quad + P(A \cap B \cap C) \end{aligned}$$

2) กฎคูณ (Multiplication Rule)

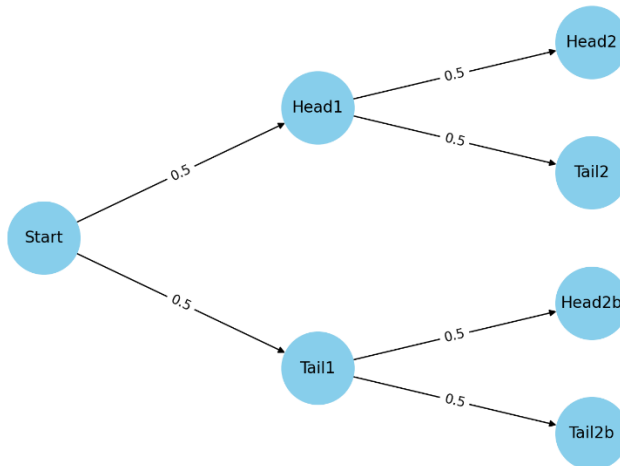
ใช้เมื่อถามความน่าจะเป็นของเหตุการณ์ “เกิดพร้อมกัน” หรือ “A และ B เกิด” (intersection)

รูปทั่วไป (ผ่านความน่าจะเป็นมีเงื่อนไข)

$$P(A \cap B) = P(A) \times P(B | A)$$

ตีความว่า โอกาสที่ “A เกิด” คูณด้วย “เมื่อ A เกิดแล้ว โอกาสที่ B จะเกิด” (Casella & Berger, 1990; สมจิต วัฒนาชาญกุล, 2546)

Multiplication Rule: Example (Tossing 2 Coins)



แผนภาพที่ 4.2 Tree Diagram

กรณีอิสระ (Independent Events)

ถ้าเหตุการณ์ไม่กระทบกัน $\Rightarrow P(B | A) = P(B)$ ดังนั้น

$$P(A \cap B) = P(A) \times P(B)$$

(Walpole, 1970)

ตัวอย่าง (เหรียญ 2 ครั้ง - อิสระ)

$$P(\text{หัวและหัว}) = 0.5 \times 0.5 = 0.25$$

ตัวอย่าง (ไฟ 52 ใบ - ไม่อิสระเมื่อ “ไม่ใส่คืน”)

1. “หัวใจทั้งสองใบ” โดยหยิบ 2 ใบต่อเนื่อง แบบไม่ใส่คืน

$$P(\text{หัวใจ, แล้วหัวใจ}) = \frac{13}{52} \times \frac{12}{51} \approx 0.0588$$

2. “เอซทั้งสองใบ” แบบไม่ใส่คืน

$$P(\text{เอซ 2 ใบ}) = \frac{4}{52} \times \frac{3}{51} = \frac{12}{2652} \approx 0.00452$$

หาก “ใส่คืนทุกครั้ง” จะเข้าเงื่อนไขอิสระ และใช้ $P(A)P(B)$ ได้ (Walpole, 1970; Milton & Arnold, 1995)

ตัวอย่างทางสังคมศาสตร์ (สำรวจแบบมีเงื่อนไข)

ให้ A = “อาศัยในเมืองใหญ่”, B = “สนับสนุนนโยบาย X”

ถ้า $P(A) = 0.40$ และจากข้อมูลพบว่าในกลุ่มคนเมืองใหญ่ 70% สนับสนุนนโยบาย X

$$P(A \cap B) = P(A) \times P(B | A) = 0.40 \times 0.70 = 0.28$$

ตีความว่า 28% ของทั้งประชากร “ทั้งอาศัยในเมืองใหญ่และสนับสนุนนโยบาย X” (Milton & Arnold, 1995)

สรุปแนวคิด-ใช้อย่างไรให้ถูก

- เมื่อคำถามเป็น “หรือ” (อย่างใดอย่างหนึ่งเกิด) $\Rightarrow$ ใช้ กฎบวก

- ถ้าเหตุการณ์ “ไม่ทับซ้อน” ให้บวกตรง ๆ

- ถ้า “ทับซ้อน” ต้องลบ $P(A \cap B)$ ออก 1 ครั้ง (อย่าลืมส่วนซ้อน)

- เมื่อคำถามเป็น “และ/พร้อมกัน” $\Rightarrow$ ใช้ กฎคูณ

- กรณีทั่วไป $P(A \cap B) = P(A) \times P(B | A)$

- ถ้า “อิสระ” $\Rightarrow$ ใช้ $P(A)P(B)$ ได้เลย

- อย่าสับสน “ไม่ทับซ้อน” และ “อิสระ” ไม่ใช่เรื่องเดียวกัน

- ถ้า ไม่ทับซ้อน (เกิดพร้อมกันไม่ได้) $\Rightarrow P(B | A) = 0$ จึงไม่อิสระ ยกเว้นกรณีเล็กน้อยที่ $P(A) = 0$ หรือ $P(B) = 0$ (Casella & Berger, 1990)

กล่าวโดยสรุป กฎทั้งสองทำให้สามารถคำนวณโอกาสของเหตุการณ์ซับซ้อนที่เกิดขึ้นร่วมกันหรือแทนกันได้อย่างเป็นระบบ และเป็นรากฐานในการวิเคราะห์ทางสถิติทุกแขนงในงานวิจัยทางสังคมศาสตร์

3. การแจกแจงแบบไม่ต่อเนื่อง (Discrete Distributions)

การแจกแจงแบบไม่ต่อเนื่อง ใช้อธิบายความน่าจะเป็นของตัวแปรสุ่มไม่ต่อเนื่อง ซึ่งมีค่าที่นับได้ เช่น จำนวนเหรียญที่ออกหัวจากการโยน 10 ครั้ง จำนวนผู้มาลงทะเบียนอบรมในแต่ละวัน หรือจำนวนครั้งที่ประชาชนเข้าร่วมกิจกรรมทางการเมือง ตัวอย่างที่ใช้กันแพร่หลายคือ การแจกแจงทวินาม (Binomial Distribution) และการแจกแจงปัวส์ซอง (Poisson Distribution) (Milton & Arnold, 1995)

1) การแจกแจงทวินาม (Binomial Distribution)

การแจกแจงทวินามเหมาะสำหรับการทดลองที่มีลักษณะดังนี้

1. การทดลองมีเพียง 2 ผลลัพธ์ (เช่น สำเร็จ/ล้มเหลว, ใช่/ไม่ใช่, เห็นด้วย/ไม่เห็นด้วย)

2. มีจำนวนการทดลองคงที่เท่ากับ n ครั้ง

3. ความน่าจะเป็นของความสำเร็จในแต่ละครั้งคงที่เท่ากับ P

4. การทดลองแต่ละครั้งเป็นอิสระต่อกัน (Walpole, 1970)

สูตรการแจกแจง

$$P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$$

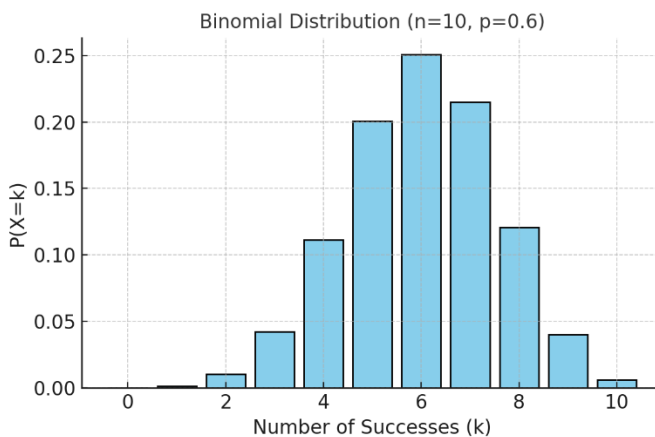
โดยที่ $k = 0, 1, 2, \dots, n$

ตัวอย่าง สมมติทำการสำรวจนักศึกษา 10 คน ว่า “เห็นด้วยกับการจัดตั้งสภานักศึกษา” หรือไม่ โดยแต่ละคนมีโอกาสเห็นด้วย 0.6 เท่ากัน และอิสระต่อกัน

- ต้องการหาความน่าจะเป็นที่มีนักศึกษา 7 คน “เห็นด้วย”

$$P(X = 7) = \binom{10}{7} (0.6)^7 (1 - 0.6)^{10-7} = 120 \times 0.02799 \times 0.064 = 0.215$$

คำตอบ ความน่าจะเป็นที่มี 7 คนเห็นด้วยคือประมาณ 21.5%



แผนภาพที่ 4.3 การแจกแจงทวินาม (Binomial Distribution)

แสดงความน่าจะเป็นของจำนวน “ความสำเร็จ” ในการทดลอง 10 ครั้ง ที่โอกาสสำเร็จในแต่ละครั้งเท่ากับ 0.6 ตัวอย่างการประยุกต์ในสังคมศาสตร์ เช่น การสำรวจความคิดเห็นของประชาชนว่ามี “ผู้เห็นด้วย” ก็คนในกลุ่มตัวอย่าง 10 คน การแจกแจงทวินามช่วยให้เราคำนวณโอกาสที่จะมีจำนวนผู้เห็นด้วยเท่ากับ k คนได้อย่างแม่นยำ

2) การแจกแจงปัวส์ซอง (Poisson Distribution)

การแจกแจงปัวส์ซองใช้กับตัวแปรสุ่มที่นับ “จำนวนเหตุการณ์” ที่เกิดขึ้นในช่วงเวลา หรือพื้นที่ที่กำหนด โดยมีเงื่อนไขว่า

1. เหตุการณ์เกิดขึ้นอย่างอิสระต่อกัน
2. ความน่าจะเป็นของการเกิดเหตุการณ์ในช่วงเวลาเล็ก ๆ สัดส่วนกับความยาวช่วงเวลา
3. ค่าเฉลี่ยของจำนวนครั้งที่เกิดในช่วงนั้นคงที่ เรียกว่า λ (Milton & Arnold, 1995)

สูตรการแจกแจง

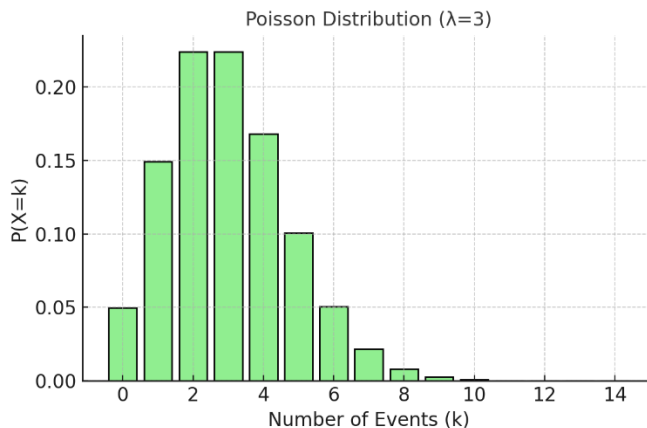
$$P(X = k) = \frac{e^{-\lambda} \lambda^k}{k!}$$

โดยที่ $k = 0, 1, 2, \dots$

ตัวอย่าง สำนักงานทะเบียนอำเภอแห่งหนึ่งมีผู้มาลงทะเบียนเกิดเฉลี่ย 3 รายต่อชั่วโมง ต้องการหาความน่าจะเป็นที่จะมี “5 ราย” มาลงทะเบียนภายใน 1 ชั่วโมง

$$P(X = 5) = \frac{e^{-3} \times 3^5}{5!} = \frac{0.0498 \times 243}{120} \approx 0.1008$$

คำตอบ ความน่าจะเป็นที่มีคนมาลงทะเบียน 5 รายใน 1 ชั่วโมงคือ ประมาณ 10.1%



แผนภาพที่ 4.4 การแจกแจงปัวส์ซอง (Poisson Distribution)

แสดงความน่าจะเป็นของจำนวนเหตุการณ์ที่เกิดขึ้นเมื่อค่าเฉลี่ย (λ) เท่ากับ 3 ตัวอย่างการประยุกต์ เช่น การวิเคราะห์จำนวนผู้มาลงทะเบียนนอกรอบในแต่ละชั่วโมง หรือจำนวนเหตุอาชญากรรมที่เกิดขึ้นในชุมชนหนึ่ง ๆ ต่อเดือน การแจกแจงปัวส์ซองเหมาะสำหรับข้อมูลที่เป็น “การนับจำนวนครั้ง” ของเหตุการณ์ที่เกิดขึ้นในช่วงเวลา/พื้นที่ที่กำหนด

สรุปได้ว่า การแจกแจงแบบไม่ต่อเนื่อง ใช้อธิบายความน่าจะเป็นของตัวแปรสุ่มไม่ต่อเนื่อง ซึ่งมีค่าที่นับได้ ตัวอย่างที่ใช้กันแพร่หลายคือ การแจกแจงทวินาม ใช้เมื่อมีจำนวนการทดลองชัดเจน แต่ครั้งมีเพียง 2 ผลลัพธ์ และความน่าจะเป็นคงที่ เช่น การโยนเหรียญ การสำรวจ “เห็นด้วย/ไม่เห็นด้วย” และการแจกแจงปัวส์ซอง ใช้เมื่อสนใจจำนวนเหตุการณ์ที่เกิดขึ้นในช่วงเวลา หรือพื้นที่ เช่น จำนวนลูกค้าเข้าร้าน จำนวนอีเมลที่เข้ามา หรือจำนวนเหตุอาชญากรรมรายเดือน

4. การแจกแจงแบบต่อเนื่อง (Continuous Distribution)

การแจกแจงแบบต่อเนื่อง ใช้อธิบายความน่าจะเป็นของตัวแปรสุ่มต่อเนื่อง ซึ่งสามารถรับค่าได้เป็นช่วงของจำนวนจริง เช่น ส่วนสูง น้ำหนัก รายได้ หรือคะแนนสอบ ตัวแปรสุ่มต่อเนื่องแตกต่างจากตัวแปรสุ่มไม่ต่อเนื่องตรงที่ ความน่าจะเป็นที่จะได้ค่าใดค่าหนึ่งพอดีมีค่าเท่ากับศูนย์ แต่จะการใช้การหาความน่าจะเป็นในลักษณะพื้นที่ใต้โค้งของฟังก์ชันความหนาแน่นความน่าจะเป็น (Probability Density Function: PDF)

1) การแจกแจงปกติ (Normal Distribution)

การแจกแจงปกติเป็นการแจกแจงแบบต่อเนื่องที่สำคัญที่สุด ใช้กันอย่างแพร่หลายในการวิเคราะห์ข้อมูลทางสังคมศาสตร์ เนื่องจากข้อมูลจำนวนมากในธรรมชาติ และสังคมมักมีลักษณะใกล้เคียงการแจกแจงปกติ เช่น คะแนนสอบ รายได้ครัวเรือน หรือระดับความพึงพอใจของประชาชน (Dudewicz, 1988; Milton & Arnold, 1995)

ลักษณะสำคัญ

- รูปร่างเป็นโค้งระฆังคว่ำ (bell-shaped curve)
- สมมาตรรอบค่าเฉลี่ย μ
- ถูกกำหนดด้วย 2 พารามิเตอร์ คือ ค่าเฉลี่ย μ และส่วนเบี่ยงเบน

มาตรฐาน σ ซึ่งบ่งบอกตำแหน่งและการกระจายของข้อมูล

สูตรฟังก์ชันความหนาแน่น

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

ตัวอย่าง สมมติคะแนนสอบวิชาสถิติของนักศึกษามีการแจกแจงแบบปกติ โดยมีค่าเฉลี่ย $\mu = 70$ และส่วนเบี่ยงเบนมาตรฐาน $\sigma = 10$ ต้องการหาความน่าจะเป็นที่นักศึกษาคนหนึ่งจะได้คะแนนมากกว่า 80

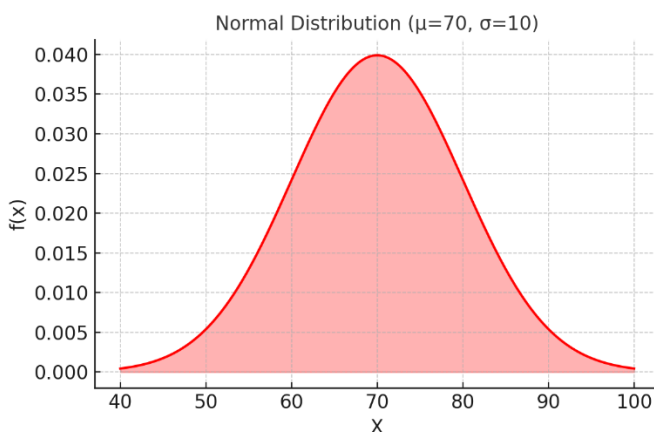
คำนวณ Z-score

$$Z = \frac{X - \mu}{\sigma} = \frac{80 - 70}{10} = 1.0$$

จากตาราง Z: $P(Z < 1.0) = 0.8413$

ดังนั้น $P(X > 80) = 1 - 0.8413 = 0.1587$

คำตอบ นักศึกษามีโอกาส 15.87% ที่จะได้คะแนนเกิน 80 คะแนน



แผนภาพที่ 4.5 การแจกแจงปกติ (Normal Distribution)

แสดงรูปโค้งการแจกแจงปกติที่มีค่าเฉลี่ย (μ) เท่ากับ 70 และส่วนเบี่ยงเบนมาตรฐาน (σ) เท่ากับ 10 ตัวอย่างการประยุกต์ในสังคมศาสตร์ เช่น การแจกแจงคะแนนสอบของนักศึกษา การวัดระดับทัศนคติทางการเมือง หรือการวิเคราะห์รายได้ประชาชน ซึ่งมักมีลักษณะใกล้เคียงการแจกแจงปกติ โค้งปกติยังเป็นรากฐานในการคำนวณค่ามาตรฐาน (Z-score) เพื่อใช้ทดสอบสมมติฐานและสร้างช่วงความเชื่อมั่น

2) สมบัติของการแจกแจงปกติ

การแจกแจงปกติ (Normal Distribution) เป็นการแจกแจงทางสถิติที่สำคัญที่สุด เนื่องจากมีคุณสมบัติทางคณิตศาสตร์ที่โดดเด่นและสามารถนำไปประยุกต์ใช้ได้กว้างขวางในเกือบทุกสาขาวิชาที่ใช้สถิติ ไม่ว่าจะเป็นวิทยาศาสตร์ เศรษฐศาสตร์ จิตวิทยา รวมถึงรัฐศาสตร์และสังคมศาสตร์ (Casella & Berger, 1990; Milton & Arnold, 1995)

สมบัติของการแจกแจงปกติช่วยให้นักวิจัยสามารถทำความเข้าใจลักษณะของข้อมูล การกระจายของตัวแปร และนำไปใช้ในการทดสอบสมมติฐานและการอนุมานเชิง

สถิติ ตัวอย่างเช่น การประเมินผลการเรียน การวิเคราะห์รายได้ประชากร หรือการวัดระดับความพึงพอใจของประชาชน สมบัติหลักที่ควรทำความเข้าใจมีดังนี้

(1) โค้งสมมาตรที่ค่าเฉลี่ย μ

การแจกแจงปกติมีลักษณะโค้งสมมาตร (symmetric) รอบค่าเฉลี่ย μ ดังนั้นโอกาสที่ข้อมูลจะอยู่ทางด้านซ้ายหรือด้านขวาของค่าเฉลี่ยจะเท่ากันพอดี

สูตร

$$P(X \leq \mu) = P(X \geq \mu) = 0.5$$

แทนค่า

ให้ $X \sim N(70, 10^2)$ หมายถึง คะแนนสอบมีค่าเฉลี่ย 70 และ $\sigma = 10$

$$P(X \leq 70) = P(X \geq 70) = 0.5$$

การนำไปใช้ ช่วยให้นักวิจัยสามารถอธิบายได้ว่ากลุ่มตัวอย่างครึ่งหนึ่งมีค่าต่ำกว่าหรือเท่ากับค่าเฉลี่ย และอีกครึ่งหนึ่งสูงกว่าหรือเท่ากับค่าเฉลี่ย เช่น ในการสำรวจความพึงพอใจ หากค่าเฉลี่ย = 3.5 (เต็ม 5) จะทราบว่า ~50% ของผู้ตอบมีคะแนนมากกว่าหรือเท่ากับ 3.5

(2) ค่าเฉลี่ย = มัธยฐาน = ฐานนิยม

เพราะโค้งสมมาตรและจุดสูงสุดอยู่ที่ค่าเฉลี่ย μ ทำให้ทั้ง ค่าเฉลี่ย (mean), มัธยฐาน (median) และ ฐานนิยม (mode) มีค่าเท่ากัน

สูตร

$$\mu = \text{median} = \text{mode}$$

แทนค่า

ให้ $X \sim N(70, 15^2) \rightarrow \mu = 100$

ดังนั้น มัธยฐานและฐานนิยมก็เท่ากับ 100 เช่นกัน

การนำไปใช้ ใช้ในงานพรรณนาสถิติ (descriptive statistics) เพราะสามารถแทนแนวโน้มเข้าสู่ศูนย์กลางด้วยค่าเดียวได้ทันที เหมาะกับการรายงานผลการสำรวจ เช่น คะแนนสอบเฉลี่ย = มัธยฐาน = ฐานนิยม = 70

(3) พื้นที่ใต้โค้งรวม = 1

โค้งปกติเป็นฟังก์ชันความหนาแน่นความน่าจะเป็น (PDF) ซึ่งมีคุณสมบัติว่าพื้นที่ใต้โค้งทั้งหมดต้องรวมกันได้ 1 (หรือ 100%)

สูตร

$$\int_{-\infty}^{\infty} f(x) dx = 1$$

โดยที่

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

แทนค่า

ถ้า $X \sim N(70, 10^2)$

$$f(x) = \frac{1}{10\sqrt{2\pi}} e^{-\frac{(x-70)^2}{200}}$$

ดังนั้น

$$\int_{-\infty}^{\infty} \frac{1}{10\sqrt{2\pi}} e^{-\frac{(x-70)^2}{200}} dx = 1$$

การนำไปใช้ ใช้เป็นหลักการพื้นฐานของการหาความน่าจะเป็นในช่วง เช่น

$$P(60 \leq X \leq 80) = \int_{60}^{80} \frac{1}{10\sqrt{2\pi}} e^{-\frac{(x-70)^2}{200}} dx$$

ในทางปฏิบัติ ใช้ตารางค่า Z หรือโปรแกรมสถิติ (เช่น SPSS, R, Python) เพื่อหาค่าพื้นที่ใต้โค้ง

(4) กฎ 68–95–99.7 (Empirical Rule)

กฎนี้อธิบายการกระจายข้อมูลภายใต้โค้งปกติ โดยกำหนดช่วงรอบค่าเฉลี่ย μ ตามจำนวนส่วนเบี่ยงเบนมาตรฐาน σ

สูตร

$$P(\mu - 1\sigma \leq X \leq \mu + 1\sigma) \approx 0.68$$

$$P(\mu - 2\sigma \leq X \leq \mu + 2\sigma) \approx 0.95$$

$$P(\mu - 3\sigma \leq X \leq \mu + 3\sigma) \approx 0.997$$

แทนค่า

สมมติ $X \sim N(70, 10^2)$

- ช่วง $\mu \pm 1\sigma = 70 \pm 10 = [60, 80] \rightarrow$ มีนักศึกษา $\sim 68\%$ ที่ได้คะแนนในช่วงนี้

- ช่วง $\mu \pm 2\sigma = [50, 90] \rightarrow \sim 95\%$ ของนักศึกษาอยู่ในช่วงนี้

- ช่วง $\mu \pm 3\sigma = [40, 100] \rightarrow \sim 99.7\%$ ของนักศึกษาอยู่ในช่วงนี้

การนำไปใช้ กฎนี้ใช้ประเมินข้อมูลเชิงประจักษ์ได้รวดเร็ว เช่น

- ในการสอบ ถ้าคะแนนเต็ม 100 ค่าเฉลี่ย = 70, $\sigma = 10 \rightarrow$ จะรู้ว่ามึนักศึกษาส่วนใหญ่ (~95%) ได้คะแนนระหว่าง 50–90

- ในงานวิจัยสังคมศาสตร์ ใช้ระบุ “กลุ่มปกติ” กับ “กลุ่มเบี่ยงเบน” เช่น การหาผู้ที่ให้คะแนนความพึงพอใจที่อยู่ห่างจากค่าเฉลี่ยมากกว่า 2 หรือ 3 ส่วนเบี่ยงเบนมาตรฐาน ถือว่าเป็น outlier

สรุปได้ว่า การแจกแจงปกติมีสมบัติสำคัญ 4 ประการ ได้แก่ โค้งสมมาตรที่ค่าเฉลี่ย ทำให้ข้อมูลกระจายสมดุลทั้งสองด้าน ค่าเฉลี่ย มัธยฐาน และฐานนิยมมีค่าเท่ากัน พื้นที่ใต้โค้งรวมเท่ากับ 1 หรือ 100% และกฎ 68–95–99.7 ที่อธิบายการกระจายข้อมูลรอบค่าเฉลี่ย สมบัติเหล่านี้ทำให้การแจกแจงปกติเป็นรากฐานของสถิติอนุมาน และถูกนำไปใช้ในการทดสอบสมมติฐาน การสร้างช่วงความเชื่อมั่น และการวิเคราะห์ข้อมูลเชิงปริมาณอย่างกว้างขวางในงานวิจัยสังคมศาสตร์ เศรษฐศาสตร์ และการเมือง

5. การใช้ตาราง Z

การแจกแจงปกติทั่วไปสามารถเปลี่ยนเป็น การแจกแจงปกติมาตรฐาน (Standard Normal Distribution) ได้ เพื่อให้ง่ายต่อการคำนวณความน่าจะเป็น เนื่องจาก การแจกแจงปกติมีหลายค่าเฉลี่ย (μ) และส่วนเบี่ยงเบนมาตรฐาน (σ) ที่แตกต่างกัน จึงต้องมีวิธีมาตรฐานในการเปรียบเทียบ ซึ่งทำได้โดยการแปลงค่าต่าง ๆ ให้อยู่ในรูปของ Z -score (Casella & Berger, 1990; Milton & Arnold, 1995; Walpole, 1970)

การคำนวณ Z -score

สูตรการคำนวณคือ

$$Z = \frac{X - \mu}{\sigma}$$

โดยที่

- X = ค่าที่สนใจ
- μ = ค่าเฉลี่ยของประชากรหรือตัวอย่าง
- σ = ส่วนเบี่ยงเบนมาตรฐาน

Z-score บอกได้ว่า ค่าของตัวแปร X อยู่ห่างจากค่าเฉลี่ยกี่ส่วนเบี่ยงเบนมาตรฐาน เช่น $Z = 1$ หมายความว่าค่าของ X สูงกว่าค่าเฉลี่ยหนึ่งส่วนเบี่ยงเบนมาตรฐาน

หลักการอ่านตาราง Z

ตาราง Z ให้ค่า พื้นที่ใต้โค้งทางด้านซ้ายของค่า Z หรือ $P(Z < z)$ ซึ่งใช้เป็นพื้นฐานในการหาความน่าจะเป็นในสถานการณ์ต่าง ๆ ได้แก่

1. หาความน่าจะเป็นระหว่างสองค่า

$$P(a < Z < b) = P(Z < b) - P(Z < a)$$

2. หาความน่าจะเป็นด้านขวาของค่า Z

$$P(Z < z) = 1 - P(Z < z)$$

3. กรณีต้องการความน่าจะเป็นด้านซ้ายของค่า Z สามารถอ่านจากตารางโดยตรง

แนวทางนี้สะดวกและแม่นยำ โดยเฉพาะเมื่อใช้ร่วมกับตารางมาตรฐานหรือซอฟต์แวร์สถิติ เช่น SPSS, R หรือ Python (Dudewicz, 1988; Milton & Arnold, 1995)

ตัวอย่างการคำนวณ

โจทย์ คะแนนสอบมีค่าเฉลี่ย $\mu = 70$, ส่วนเบี่ยงเบนมาตรฐาน $\sigma = 10$
ต้องการหาความน่าจะเป็นที่นักเรียนจะได้คะแนนมากกว่า 80

1. แปลงเป็น Z -score

$$Z = \frac{X - \mu}{\sigma} = \frac{80 - 70}{10} = 1.0$$

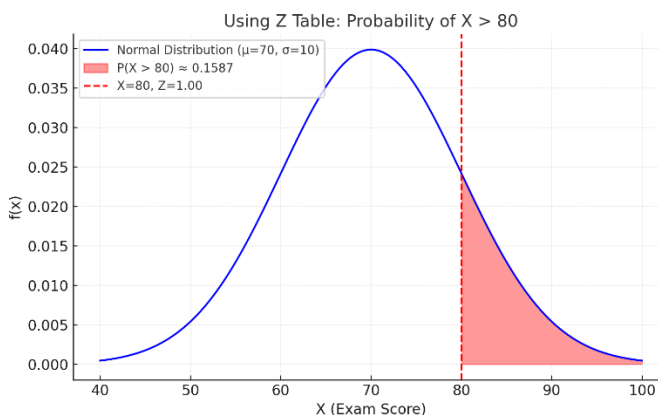
2. หาจากตาราง Z

$$P(Z < 1.0) = 0.8413$$

3. หาความน่าจะเป็นด้านขวา

$$P(Z < 1.0) = 1 - 0.8413 = 0.1587$$

คำตอบ โอกาสที่นักเรียนจะได้คะแนนมากกว่า 80 คือ 15.87%



แผนภาพที่ 4.6 แสดงการใช้ตาราง Z

ในการหาความน่าจะเป็นที่นักเรียนได้คะแนนมากกว่า 80 ภายใต้การแจกแจงปกติที่มีค่าเฉลี่ย $\mu = 70$ และส่วนเบี่ยงเบนมาตรฐาน $\sigma = 10$ โดยทำการแปลงคะแนนเป็น Z -score ได้ $Z = 1.0$ พื้นที่ใต้โค้งทางด้านขวาของเส้นแนวตั้งที่ $X = 80$ คือ 0.1587

หรือร้อยละ 15.87 ของพื้นที่ทั้งหมด ซึ่งหมายความว่ามีความน่าจะเป็น 15.87% ที่นักเรียนจะได้คะแนนมากกว่า 80

การนำไปใช้ในงานวิจัยทางสังคมศาสตร์

1. การศึกษา ใช้ประเมินผลสัมฤทธิ์ของนักเรียนโดยเปรียบเทียบว่าใครอยู่สูงหรือต่ำกว่าค่าเฉลี่ยที่ส่วนเบี่ยงเบนมาตรฐาน เช่น นักเรียนที่มี $Z = 2$ แสดงว่ามีคะแนนสูงกว่าค่าเฉลี่ยถึง 2 เท่าของ σ ซึ่งถือว่าสูงมาก

2. การสำรวจความคิดเห็น ใช้ตรวจสอบว่าคะแนนความพึงพอใจของกลุ่มใด ๆ สูงหรือต่ำกว่ากลุ่มประชากรหลักอย่างมีนัยสำคัญหรือไม่

3. เศรษฐศาสตร์และการเมือง ใช้วิเคราะห์ความเสี่ยงและการคาดการณ์ เช่น ความน่าจะเป็นที่อัตราการว่างงานจะเกินค่าหนึ่ง หรือโอกาสที่อัตราเงินเฟ้อจะต่ำกว่าระดับที่กำหนด

สรุปได้ว่า การใช้ตาราง Z เป็นขั้นตอนสำคัญในการวิเคราะห์ข้อมูลที่มีการแจกแจงใกล้เคียงปกติ โดยช่วยให้ผู้วิจัยสามารถตีความความน่าจะเป็นในช่วงต่าง ๆ ของข้อมูลได้อย่างแม่นยำ ทำให้สามารถอธิบายปรากฏการณ์ทางสังคมศาสตร์ การศึกษา และเศรษฐศาสตร์ได้ดียิ่งขึ้น

6. ความสำคัญของความน่าจะเป็นต่อการวิเคราะห์เชิงอนุมาน

ความน่าจะเป็น (Probability) ถือเป็นรากฐานสำคัญของ สถิติอนุมาน (Inferential Statistics) เนื่องจากช่วยให้นักวิจัยสามารถอธิบายและตัดสินใจภายใต้สภาวะที่มีความไม่แน่นอน (Casella & Berger, 1990; Walpole, 1970) กล่าวคือ ในการวิจัยทางสังคมศาสตร์ นักวิจัยไม่สามารถศึกษาประชากรทั้งหมดได้ แต่ใช้การสุ่มตัวอย่างเพื่อนำมาสรุปอธิบายประชากร ดังนั้นความน่าจะเป็นจึงทำหน้าที่เป็น “สะพานเชื่อม” ระหว่างข้อมูลที่ได้จากตัวอย่างกับข้อสรุปที่เกี่ยวข้องกับประชากร

1) ใช้สร้างแบบจำลองทางสถิติ

แบบจำลองทางสถิติ (Statistical Models) อาศัยแนวคิดของความน่าจะเป็น เช่น การกำหนดว่า “คะแนนสอบ” อาจมีการแจกแจงแบบปกติ หรือ “จำนวนผู้เข้าร่วมกิจกรรม” อาจมีการแจกแจงแบบปัวซอง ซึ่งช่วยให้นักวิจัยสามารถเลือกแบบจำลองที่เหมาะสมในการอธิบายข้อมูลจริง (Milton & Arnold, 1995) การสร้างแบบจำลองเช่นนี้ช่วยให้สามารถทำนายผลลัพธ์ในอนาคตและจำลองสถานการณ์ทางสังคมที่ซับซ้อนได้

2) เป็นพื้นฐานของการทดสอบสมมติฐาน (Hypothesis Testing)

การทดสอบสมมติฐานต้องอาศัยความน่าจะเป็นในการกำหนด ค่าความน่าจะเป็นที่ผิดพลาด (p -value) เพื่อตัดสินใจว่าจะ “ปฏิเสธ” หรือ “ยอมรับ” สมมติฐานศูนย์

ตัวอย่างเช่น ในการสำรวจความคิดเห็น หากต้องการตรวจสอบว่าประชาชน “ส่วนใหญ่” สนับสนุนนโยบายใหม่หรือไม่ จะต้องคำนวณค่า p -value และเปรียบเทียบกับระดับนัยสำคัญที่กำหนด เช่น 0.05 (Casella & Berger, 1990)

3) ใช้ในการสร้างช่วงความเชื่อมั่น (Confidence Interval)

ความน่าจะเป็นเป็นพื้นฐานในการกำหนดช่วงความเชื่อมั่น เช่น ช่วงความเชื่อมั่น 95% หมายถึง ถ้าสุ่มตัวอย่างซ้ำหลายครั้ง จะได้ค่าประมาณที่ครอบคลุมค่าเฉลี่ยจริงของประชากรประมาณ 95% (Walpole, 1970) สิ่งนี้ช่วยให้นักวิจัยสามารถรายงานผลวิจัยพร้อมระดับความเชื่อมั่นที่ชัดเจน ไม่ใช่เพียงค่าเฉลี่ยเพียงจุดเดียว

4) ใช้ในการประเมินความเสี่ยงและความน่าเชื่อถือของผลวิจัย

ความน่าจะเป็นทำให้เราสามารถประเมิน “ความเสี่ยงที่จะสรุปผิด” ได้ เช่น การยอมรับผลการทดสอบสมมติฐานทั้งที่ไม่จริง (Type I Error) หรือการไม่ยอมรับทั้งที่จริง (Type II Error) (Dudewicz, 1988) นอกจากนี้ยังช่วยให้นักวิจัยสามารถประเมินความน่าเชื่อถือ (reliability) ของผลลัพธ์ได้ ซึ่งมีความสำคัญอย่างยิ่งต่อการตัดสินใจเชิงนโยบายและการบริหารจัดการสังคม

5) ช่วยในการตัดสินใจบนฐานข้อมูล (Data-Driven Decision Making)

การใช้ความน่าจะเป็นทำให้การตัดสินใจของนักวิจัยและผู้บริหารไม่ขึ้นอยู่กับการคาดเดา แต่ตั้งอยู่บนพื้นฐานของข้อมูลและการวิเคราะห์เชิงปริมาณ ตัวอย่างเช่น การพิจารณาว่า “นโยบายใหม่มีแนวโน้มประสบความสำเร็จ” สามารถใช้แบบจำลองความน่าจะเป็นเพื่อประเมินโอกาสและลดความเสี่ยงได้ (Milton & Arnold, 1995)

สรุปได้ว่า ความน่าจะเป็นจึงไม่ใช่เพียงเครื่องมือทางคณิตศาสตร์ แต่เป็นรากฐานของการวิเคราะห์เชิงอนุมานที่ช่วยให้นักวิจัยสามารถ สร้างแบบจำลอง ทดสอบสมมติฐาน สร้างช่วงความเชื่อมั่น ประเมินความเสี่ยง และ ตัดสินใจอย่างมีข้อมูลรองรับสมบัติเหล่านี้ทำให้ความน่าจะเป็นเป็นหัวใจของการวิจัยเชิงปริมาณและการตัดสินใจเชิงนโยบายในงานสังคมศาสตร์ เศรษฐศาสตร์ และรัฐศาสตร์

สรุปท้ายบท

ความน่าจะเป็นและการแจกแจงความน่าจะเป็นเป็นหัวใจสำคัญของการศึกษาทางสถิติและการวิจัยทางสังคมศาสตร์ ทั้งนี้เพราะความน่าจะเป็นทำหน้าที่เป็นเครื่องมือในการอธิบายความไม่แน่นอนของปรากฏการณ์ต่าง ๆ ได้อย่างมีระบบและมีเหตุผล การทำความเข้าใจแนวคิดพื้นฐานของความน่าจะเป็นช่วยให้นักวิจัยสามารถตีความข้อมูลจากการสุ่มตัวอย่างและขยายข้อสรุปไปสู่ประชากรได้อย่างน่าเชื่อถือ

ในบทนี้ได้อธิบายถึงความหมายและแนวคิดของความน่าจะเป็น ซึ่งครอบคลุมการตีความทั้งเชิงทฤษฎี เชิงประจักษ์ และเชิงอัตนัย ตลอดจนกฎของความน่าจะเป็นที่ใช้เป็นเครื่องมือในการคำนวณโอกาสของเหตุการณ์ซับซ้อน ไม่ว่าจะเป็นการรวมเหตุการณ์ การเกิดพร้อมกัน หรือการใช้กฎเสริมเพื่อหาความน่าจะเป็นในมุมมองที่แตกต่างกัน นอกจากนี้ยังได้ศึกษาการแจกแจงความน่าจะเป็นทั้งแบบไม่ต่อเนื่อง เช่น การแจกแจงทวินามและการแจกแจงปัวส์ซอง ที่เหมาะกับข้อมูลที่เป็นการนับจำนวน และการแจกแจงแบบต่อเนื่อง โดยเฉพาะการแจกแจงปกติซึ่งถือเป็นรากฐานสำคัญของการวิเคราะห์เชิงอนุมาน

ในส่วนของการใช้ตาราง Z ผู้เรียนได้เรียนรู้การแปลงค่ามาตรฐานเป็น Z - score เพื่อคำนวณพื้นที่ใต้โค้งปกติมาตรฐาน ซึ่งช่วยให้สามารถหาความน่าจะเป็นของค่าหรือช่วงค่าที่สนใจได้อย่างสะดวกและแม่นยำ สุดท้าย บทนี้ยังได้ชี้ให้เห็นถึงความสำคัญของความน่าจะเป็นต่อการวิเคราะห์เชิงอนุมาน ซึ่งปรากฏทั้งในการสร้างแบบจำลองทางสถิติ การทดสอบสมมติฐาน การสร้างช่วงความเชื่อมั่น การประเมินความเสี่ยงของการตัดสินใจ และการทำให้การวิจัยเชิงสังคมศาสตร์มีความถูกต้องและน่าเชื่อถือมากยิ่งขึ้น

กล่าวโดยสรุป ความน่าจะเป็นและการแจกแจงความน่าจะเป็นไม่เพียงแต่เป็นทฤษฎีทางคณิตศาสตร์ หากแต่ยังเป็นภาษาสำคัญในการสื่อสารความไม่แน่นอนและสนับสนุนการตัดสินใจเชิงวิชาการและเชิงนโยบาย การเข้าใจและสามารถประยุกต์ใช้แนวคิดเหล่านี้ได้อย่างถูกต้องจึงเป็นทักษะที่จำเป็นสำหรับนักวิจัย นักศึกษา และผู้ปฏิบัติงานในสาขาสังคมศาสตร์และศาสตร์ที่เกี่ยวข้อง

บทที่ 5

การอนุมานทางสถิติ (STATISTICAL INFERENCE)

การวิจัยทางสังคมศาสตร์เป็นศาสตร์ที่ต้องเผชิญกับความซับซ้อนของพฤติกรรมมนุษย์และโครงสร้างทางสังคม ข้อมูลที่นักวิจัยเก็บได้จากการสำรวจหรือการทดลองมักมีข้อจำกัดทั้งด้านจำนวนและเวลา การใช้วิธีการทางสถิติอนุมาน (statistical inference) จึงเป็นหัวใจสำคัญที่ช่วยให้สามารถนำข้อมูลจาก ตัวอย่าง (sample) มาสรุปเป็นข้อเท็จจริงของ ประชากร (population) ได้อย่างมีเหตุผลและน่าเชื่อถือ โดยอาศัยพื้นฐานทางความน่าจะเป็นเป็นเครื่องมือหลัก (Cox, 2006)

การอนุมานทางสถิติเป็นกลไกที่ช่วยให้นักวิจัยไม่เพียงหยุดอยู่กับการบรรยายข้อมูลที่มีอยู่ แต่สามารถ “ก้าวข้าม” ข้อมูลที่จำกัด ไปสู่การสร้างข้อสรุปที่ครอบคลุมและเป็นประโยชน์ต่อการทำความเข้าใจปรากฏการณ์ทางสังคมได้กว้างขึ้น ตัวอย่างเช่น นักสังคมวิทยาสามารถใช้การทดสอบสมมติฐานเพื่อยืนยันว่าความแตกต่างของรายได้ระหว่างเพศชายและหญิงในชุมชนหนึ่งมีนัยสำคัญทางสถิติหรือไม่ หรือนักรัฐศาสตร์สามารถประมาณสัดส่วนของประชาชนที่สนับสนุนนโยบายใดนโยบายหนึ่งโดยอาศัยข้อมูลเพียงจากกลุ่มตัวอย่างที่สุ่มมาได้ (Frost, 2021)

ความสำคัญของการอนุมานทางสถิติในงานวิจัยทางสังคมศาสตร์จึงประกอบด้วยหลายประการ ได้แก่ การช่วยให้การตัดสินใจทางวิชาการตั้งอยู่บนฐานข้อมูลเชิงประจักษ์ การจัดการกับความไม่แน่นอนของข้อมูล การสนับสนุนการสร้างองค์ความรู้ใหม่ และการนำไปประยุกต์ใช้ในเชิงนโยบายหรือการแก้ปัญหาสังคม (Penn State University, 2023) ด้วยเหตุนี้ การทำความเข้าใจแนวคิดและวิธีการต่าง ๆ ของการอนุมานจึงเป็นทักษะที่สำคัญสำหรับนักวิจัย

เพื่อให้การศึกษามีความเป็นระบบ บทนี้จะนำเสนอเนื้อหาเกี่ยวกับองค์ประกอบหลักของการอนุมานทางสถิติ ไล่เรียงตั้งแต่ ประชากรและตัวอย่าง ซึ่งเป็นรากฐานของการวิเคราะห์ การประมาณค่าเฉลี่ยและสัดส่วนของประชากร และการสร้างช่วงความเชื่อมั่น (confidence interval) ซึ่งใช้เพื่อบอกขอบเขตของค่าพารามิเตอร์ที่เป็นไปได้ จากนั้นจะกล่าวถึง การทดสอบสมมติฐาน (hypothesis testing) ทั้งในเชิงหลักการทั่วไปและขั้นตอนปฏิบัติจริง ตลอดจนการอธิบายเรื่อง ค่าพี (p-value) และระดับนัยสำคัญ ซึ่งเป็นเกณฑ์สำคัญในการตีความผลลัพธ์ นอกจากนี้ยังจะมีการนำเสนอ การ

ทดสอบค่าเฉลี่ย ผ่านเครื่องมือที่ใช้กันแพร่หลายอย่าง Z-test และ t-test (One Sample) พร้อมทั้งการอธิบาย ความผิดพลาดแบบที่ 1 และแบบที่ 2 เพื่อให้ผู้อ่านตระหนักถึงข้อจำกัดและความเสี่ยงในการสรุปผลการวิจัย เนื้อหาทั้งหมดนี้จะช่วยให้เข้าใจทั้งมิติทางทฤษฎีและการประยุกต์ใช้ในงานวิจัยทางสังคมศาสตร์ได้อย่างครอบคลุม

1. ประชากรและตัวอย่าง (Population vs Sample)

การกำหนดประชากรและตัวอย่างเป็นขั้นตอนแรกที่สำคัญในการทำวิจัยเชิงปริมาณ เพราะนักวิจัยไม่สามารถเก็บข้อมูลจากประชากรทั้งหมดได้ การเลือกตัวอย่างจึงมีบทบาทสำคัญในการทำให้การวิจัยมีความเป็นไปได้ ทั้งยังต้องมีความเที่ยงตรง (validity) และสามารถอนุมานผลสู่ประชากรได้อย่างน่าเชื่อถือ

1.1 ความหมายของประชากร (Population)

ประชากร (Population) หมายถึง กลุ่มหน่วยทั้งหมดที่นักวิจัยต้องการศึกษา ไม่ว่าจะเป็นคน สิ่งของ เหตุการณ์ หรือหน่วยข้อมูลใด ๆ ที่อยู่ในขอบเขตการวิจัย ตัวอย่างเช่น

- ประชากรนักเรียนมัธยมปลายในจังหวัดนนทบุรี
- ประชากรผู้มีสิทธิเลือกตั้งในการเลือกตั้งทั่วไป
- ครั้วเรือนในชุมชนเมืองและชนบท

ประชากรมีลักษณะสำคัญคือ มีจำนวนมากและหลากหลาย การเก็บข้อมูลจากประชากรทั้งหมดมักทำไม่ได้เพราะใช้เวลาและทรัพยากรสูง (Creswell & Creswell, 2018)

1.2 ความหมายของตัวอย่าง (Sample)

ตัวอย่าง (Sample) คือ ส่วนหนึ่งของประชากรที่ถูกเลือกมาเพื่อเก็บข้อมูลและนำมาวิเคราะห์แทนประชากรทั้งหมด การเลือกตัวอย่างที่ดีต้องสะท้อนลักษณะของประชากร เพื่อให้ผลการวิจัยสามารถอนุมานได้อย่างถูกต้อง (Frost, 2021) ตัวอย่างเช่น หากนักวิจัยต้องการศึกษาพฤติกรรมการใช้สื่อสังคมออนไลน์ของนักเรียนมัธยม 10,000 คน อาจสุ่มตัวอย่าง 400 คนมาเก็บข้อมูล ซึ่งถือเป็นขนาดตัวอย่างที่เหมาะสมตามสูตรการหาขนาดตัวอย่างของ Yamane (1967)

1.3 พารามิเตอร์ (Parameter) และสถิติ (Statistic)

1) พารามิเตอร์ (Parameter) คือ ค่าที่แทนลักษณะของประชากร เช่น

- ค่าเฉลี่ยประชากร (μ) คือ ค่ากลางที่แทนระดับโดยเฉลี่ยของทุกหน่วยในประชากร เป็นตัวแทน “แนวโน้มเข้าสู่ศูนย์กลาง” ของข้อมูลทั้งหมด (Creswell & Creswell, 2018)

สูตร

$$\mu = \frac{\sum X_i}{N}$$

โดยที่

X_i = ค่าของข้อมูลแต่ละหน่วย

μ = ค่าเฉลี่ยประชากร

N = จำนวนหน่วยทั้งหมดในประชากร

ตัวอย่าง นักวิจัยเก็บข้อมูลรายได้ของเยาวชน 5 คน ได้แก่ 200, 220, 180, 210, 190

$$\mu = \frac{200+220+180+210+190}{5}$$

$$\mu = \frac{1000}{5} = 200$$

ดังนั้น ค่าเฉลี่ยรายได้ของตัวอย่างคือ 200 บาท/วัน

- ส่วนเบี่ยงเบนมาตรฐานประชากร (σ) คือ ค่าที่บอกลถึงการกระจายตัวของข้อมูลในประชากรว่ามีความห่างจากค่าเฉลี่ยมากน้อยเพียงใด ยิ่งค่า σ สูง ข้อมูลยิ่งกระจายตัวกว้าง ยิ่งค่า σ ต่ำ ข้อมูลจะรวมอยู่ใกล้ค่าเฉลี่ย (Upton & Cook, 2008)

สูตร

$$\sigma = \sqrt{\frac{\sum (X_i - \mu)^2}{N}}$$

โดยที่

X_i = ค่าของข้อมูลแต่ละหน่วย

μ = ค่าเฉลี่ยประชากร

N = จำนวนหน่วยทั้งหมดในประชากร

ตัวอย่าง ใช้ข้อมูลรายได้ 5 คน เช่นเดิม 200, 220, 180, 210, 190

- ค่าเฉลี่ยประชากร (μ) = 200

คำนวณ $(X_i - \mu)^2$:

$$- (200 - 200)^2 = 0$$

$$- (220 - 200)^2 = 400$$

$$- (180 - 200)^2 = 400$$

$$- (210 - 200)^2 = 100$$

$$- (190 - 200)^2 = 100$$

รวมผลต่างยกกำลังสอง = $0 + 400 + 400 + 100 + 100 = 1,000$

$$\sigma = \sqrt{\frac{1000}{5}} = \sqrt{200} \approx 14.14$$

ดังนั้น ส่วนเบี่ยงเบนมาตรฐานประชากร (σ) ≈ 14.14 บาท

2) สถิติ (Statistic) คือ ค่าที่คำนวณได้จากตัวอย่าง เช่น

- ค่าเฉลี่ยตัวอย่าง ($\bar{X}$) ค่าที่ได้จากการเฉลี่ยข้อมูลของกลุ่มตัวอย่างที่เก็บมาใช้ในการวิจัย ใช้แทนค่าเฉลี่ยประชากร (μ) ในการประมาณหรืออนุมานค่าพารามิเตอร์ (Frost, 2021)

สูตร

$$\bar{X} = \frac{\sum X_i}{n}$$

โดยที่

X_i = ค่าของข้อมูลแต่ละหน่วยในตัวอย่าง

$(\bar{X})$ = ค่าเฉลี่ยตัวอย่าง

n = จำนวนหน่วยทั้งหมดในตัวอย่าง

ตัวอย่าง สมมติสุ่มตัวอย่างนักเรียน 4 คน เพื่อดูจำนวนชั่วโมงการอ่านหนังสือ/วัน
ได้ข้อมูล: 2, 3, 4, 1

$$\bar{X} = \frac{2+3+4+1}{4}$$

$$\bar{X} = \frac{10}{4} = 2.5$$

ดังนั้น ค่าเฉลี่ยตัวอย่าง $\bar{X} = 2.5$ ชั่วโมง/วัน

- ส่วนเบี่ยงเบนมาตรฐานตัวอย่าง (s) คือ ค่าที่บอกการกระจายของข้อมูลในตัวอย่าง ใช้แทนค่าประชากร (σ) เพื่อทำการอนุมาน โดยสูตรการคำนวณมีการหารด้วย $n - 1$ แทนที่จะเป็น n เพื่อแก้ไขอคติในการประมาณค่า (bias correction) (Creswell & Creswell, 2018)

สูตร

$$s = \sqrt{\frac{\sum (X_i - \bar{X})^2}{n - 1}}$$

โดยที่

X_i = ค่าของข้อมูลแต่ละหน่วยในตัวอย่าง

$(\bar{X})$ = ค่าเฉลี่ยตัวอย่าง

n = จำนวนหน่วยทั้งหมดในตัวอย่าง

ตัวอย่าง ใช้ข้อมูลจำนวนชั่วโมงการอ่านหนังสือของนักเรียน 4 คน: 2, 3, 4, 1

- ค่าเฉลี่ยตัวอย่าง $\bar{X} = 2.5$

คำนวณ $(X_i - \bar{X})^2$:

$$- (2 - 2.5)^2 = 0.25$$

$$- (3 - 2.5)^2 = 0.25$$

$$- (4 - 2.5)^2 = 2.25$$

$$- (1 - 2.5)^2 = 2.25$$

ผลรวม = 5

$$s = \sqrt{\frac{5}{4-1}} = \sqrt{\frac{5}{3}} = \sqrt{1.67} \approx 1.29$$

ดังนั้น ส่วนเบี่ยงเบนมาตรฐานตัวอย่าง (s) ≈ 1.29 ชั่วโมง/วัน

1.4 ประเภทของการสุ่มตัวอย่าง

การเลือกตัวอย่าง (Sampling) เป็นกระบวนการสำคัญในการวิจัยเชิงปริมาณ เพราะผลการวิจัยที่ได้จากตัวอย่างจะถูกนำไป อนุมานแทนประชากร ทั้งหมด การเลือกวิธีสุ่มตัวอย่างที่เหมาะสมจะช่วยให้ผลลัพธ์มีความน่าเชื่อถือและลดความเอนเอียง (bias)

1) การสุ่มตัวอย่างแบบความน่าจะเป็น (Probability Sampling)

การสุ่มตัวอย่างที่ทุกหน่วยในประชากรมีโอกาสถูกเลือกเข้ามาอยู่ในกลุ่มตัวอย่างเท่ากันหรือตามความน่าจะเป็นที่กำหนดไว้ล่วงหน้า ทำให้ผลที่ได้มีความเที่ยงตรง และสามารถอ้างอิงถึงประชากรได้ (Creswell & Creswell, 2018)

(1) การสุ่มอย่างง่าย (Simple Random Sampling)

- ทุกหน่วยในประชากรมีโอกาสถูกเลือกเข้ามาเท่ากัน
- มักใช้วิธีจับสลาก รายชื่อ หรือใช้โปรแกรมคอมพิวเตอร์สุ่ม
- สูตรความน่าจะเป็นของการถูกเลือก

$$P = \frac{n}{N}$$

โดยที่ n = ขนาดตัวอย่าง, N = ขนาดประชากร

ตัวอย่าง ประชากรนักเรียน 1,000 คน เลือกตัวอย่าง 100 คน

$$P = \frac{100}{1000} = 0.10$$

ดังนั้น นักเรียนแต่ละคนมีโอกาสถูกเลือก = 10%

(2) การสุ่มแบบแบ่งชั้น (Stratified Sampling)

- ประชากรถูกแบ่งออกเป็น “ชั้น” (Strata) ตามลักษณะ เช่น เพศ อายุ ชั้นปี
- จากนั้นสุ่มตัวอย่างในแต่ละชั้นตามสัดส่วนที่เหมาะสม

- ตัวอย่าง แบ่งนักเรียนออกเป็นชายและหญิง จากนั้นสุ่มตัวอย่างตามสัดส่วนเพศที่มีอยู่จริงในประชากร

(3) การสุ่มแบบกลุ่ม (Cluster Sampling)

- เลือกทั้ง “กลุ่ม” เข้ามาเป็นตัวอย่างแทนที่จะเลือกหน่วยรายบุคคล
- ใช้เมื่อประชากรมีการกระจายกว้างและเข้าถึงยาก
- ตัวอย่าง เลือกห้องเรียนทั้งห้อง 5 ห้องจากทั้งหมด 50 ห้องในโรงเรียน

(4) การสุ่มแบบเป็นระบบ (Systematic Sampling)

- เลือกหน่วยตัวอย่างตามลำดับ เช่น ทุก ๆ คนที่ 10 จากรายชื่อทั้งหมด
- สูตรหา “ช่วงการสุ่ม (k)”

$$k = \frac{N}{n}$$

โดยที่ N = ขนาดประชากร, n = ขนาดตัวอย่าง

ตัวอย่าง มีนักศึกษา 1,200 คน ต้องการสุ่ม 120 คน

$$k = \frac{1200}{120} = 10$$

ดังนั้น เลือกทุก ๆ คนที่ 10 จากรายชื่อ → ได้ตัวอย่างครบ 120 คน

2) การสุ่มตัวอย่างแบบไม่ใช้ความน่าจะเป็น (Non-Probability Sampling)

การสุ่มตัวอย่างที่ไม่ได้อาศัยหลักความน่าจะเป็น หน่วยในประชากร ไม่ได้มีโอกาสถูกเลือกเท่ากัน มักเลือกตามความสะดวกหรือดุลยพินิจ เหมาะสำหรับการวิจัยเชิงคุณภาพหรือกรณีที่เข้าถึงประชากรได้ยาก (Etikan, Musa, & Alkassim, 2016)

(1) Convenience Sampling (การเลือกแบบสะดวก)

- เลือกตัวอย่างที่เข้าถึงง่ายที่สุด
- ตัวอย่าง แจกแบบสอบถามให้นักศึกษาที่อยู่ในห้องเรียนขณะนั้น

(2) Purposive Sampling (การเลือกแบบเจาะจง)

- เลือกตัวอย่างที่นักวิจัยเห็นว่าเหมาะสมที่สุดต่อวัตถุประสงค์
- ตัวอย่าง เลือกผู้บริหารโรงเรียนเพื่อศึกษาภาวะผู้นำทางการศึกษา

(3) Snowball Sampling (การเลือกแบบลูกโซ่)

- ใช้เครือข่ายผู้ให้ข้อมูล โดยเริ่มจากผู้หนึ่งแนะนำต่อไปยังผู้อื่น
- ตัวอย่าง ศึกษากลุ่มนักกิจกรรมทางสังคม โดยเริ่มจากผู้ให้ข้อมูล 1 คน

แล้วขยายไปยังเพื่อนในเครือข่าย

สรุปได้ว่า ประชากรคือตัวแทนทั้งหมดของสิ่งที่ศึกษา ส่วนตัวอย่างคือหน่วยย่อยที่นำมาวิเคราะห์แทนประชากร การทำความเข้าใจพารามิเตอร์และสถิติ รวมถึงการ

เลือกวิธีสุ่มตัวอย่างที่เหมาะสม เป็นหัวใจสำคัญที่ทำให้การวิจัยเชิงปริมาณมีความน่าเชื่อถือ สามารถอนุมานผลลัพธ์ไปยังประชากรได้อย่างมีประสิทธิภาพและแม่นยำ

2. การประมาณค่าเฉลี่ยและสัดส่วนของประชากร

การวิจัยเชิงปริมาณมีเป้าหมายสำคัญคือการใช้ข้อมูลจากตัวอย่าง (sample) เพื่อนำไปอนุมานถึงประชากร (population) ทั้งหมด อย่างไรก็ตาม การใช้ตัวอย่างแทนประชากรย่อมมีความคลาดเคลื่อน ดังนั้นจึงต้องมีวิธีการทางสถิติที่ช่วยให้นักวิจัยสามารถประมาณค่าของพารามิเตอร์ประชากรได้อย่างมีระบบ กระบวนการนี้เรียกว่า การประมาณค่า (Estimation) ซึ่งมีทั้งการประมาณค่าแบบจุด (point estimation) และการประมาณค่าแบบช่วง

2.1 การประมาณค่า

การประมาณค่า (Estimation) เป็นกระบวนการสำคัญในการวิจัยเชิงปริมาณ เนื่องจากนักวิจัยไม่สามารถเก็บข้อมูลจากทุกหน่วยในประชากรได้ จึงต้องอาศัยข้อมูลจากตัวอย่าง (sample) เพื่อนำมาสรุปหรืออนุมานแทนค่าพารามิเตอร์ของประชากร (population parameter) เช่น ค่าเฉลี่ยประชากร (μ) หรือสัดส่วนประชากร (p)

การประมาณค่ามี 2 รูปแบบหลัก (Cox, 2006; Penn State University, 2023) ดังนี้

1) การประมาณค่าแบบจุด (Point Estimation)

ใช้ค่าสถิติจากตัวอย่างเป็นค่าประมาณเดียว เช่น ค่าเฉลี่ยตัวอย่าง ($\bar{X}$) ใช้แทน μ หรือ $\hat{p}$ ใช้แทน p

2) การประมาณค่าแบบช่วง (Interval Estimation)

ใช้ค่าสถิติจากตัวอย่างร่วมกับทฤษฎีการแจกแจงและค่าความเชื่อมั่น เพื่อสร้างช่วงความเชื่อมั่น (Confidence Interval: CI) ที่บอกว่า พารามิเตอร์จริงของประชากรมีโอกาสอยู่ในช่วงนั้น เช่น ช่วงความเชื่อมั่น 95%

การประมาณค่าจึงช่วยให้นักวิจัยจัดการกับความไม่แน่นอน (uncertainty) ของข้อมูล และสามารถตัดสินใจได้อย่างเป็นระบบ (Frost, 2021)

2.2 การประมาณค่าเฉลี่ย (Mean Estimation)

1) การประมาณค่าแบบจุด (Point Estimation)

สูตร

$$\bar{x} = \frac{\sum X_i}{n}$$

โดยที่

X_i = ค่าของข้อมูลแต่ละหน่วยในตัวอย่าง

$(\bar{X})$ = ค่าเฉลี่ยตัวอย่าง

n = จำนวนหน่วยทั้งหมดในตัวอย่าง

ตัวอย่าง นักวิจัยต้องการศึกษาคะแนนสอบสถิติของนักศึกษา 5 คน ได้แก่ 70, 65, 80, 75, 70

$$\bar{X} = \frac{70+65+80+75+70}{5}$$

$$\bar{X} = \frac{360}{5} = 72$$

ดังนั้น ค่าเฉลี่ยตัวอย่าง $\bar{X} = 72$ คะแนน ใช้แทนค่าเฉลี่ยประชากร (μ)

(2) การประมาณค่าแบบช่วง (Interval Estimation)

การสร้างช่วงความเชื่อมั่น (Confidence Interval: CI) ช่วยให้นักวิจัยไม่เพียงแต่รู้ค่าเฉลี่ยของตัวอย่าง แต่ยังเข้าใจขอบเขตที่ค่าเฉลี่ยประชากรมีโอกาสอยู่จริง

- กรณีทราบค่า σ (ใช้ Z-distribution)

$$CI = \bar{X} \pm Z_{\alpha/2} \cdot \frac{\sigma}{\sqrt{n}}$$

- กรณีไม่ทราบค่า σ (ใช้ s และ t-distribution)

$$CI = \bar{X} \pm t_{\alpha/2, n-1} \cdot \frac{s}{\sqrt{n}}$$

ตัวอย่าง นักวิจัยต้องการประมาณค่าเฉลี่ยรายได้ต่อวันของนักศึกษา

- ค่าเฉลี่ยตัวอย่าง ($\bar{X}$) = 500 บาท
- ส่วนเบี่ยงเบนมาตรฐานตัวอย่าง (s) = 50 บาท
- ขนาดตัวอย่าง (n) = 25
- ระดับความเชื่อมั่น 95% ($t_{0.025, 24} \approx 2.064$)

$$CI = 500 \pm 2.064 \cdot \frac{50}{\sqrt{25}} = 500 \pm 2.064 \cdot 10 = 500 \pm 20.64$$

ดังนั้น ช่วงความเชื่อมั่น 95% = (479.36, 520.64) บาท

การตีความ เรามั่นใจ 95% ว่าค่าเฉลี่ยรายได้ประชากรจริงจะอยู่ระหว่าง 479.36–520.64 บาทต่อวัน

2.3 การประมาณค่าสัดส่วน (Proportion Estimation)

(1) การประมาณค่าแบบจุด (Point Estimation)

สูตร

$$\hat{p} = \frac{x}{n}$$

โดยที่

x = จำนวนหน่วยที่มีลักษณะตรงตามเกณฑ์

n = ขนาดตัวอย่าง

ตัวอย่าง จากการสำรวจนักเรียน 200 คน พบว่า 120 คนใช้สมาร์ทโฟนเพื่อการเรียน

$$\hat{p} = \frac{120}{200} = 0.60$$

ดังนั้น สัดส่วนของนักเรียนที่ใช้สมาร์ทโฟนเพื่อการเรียน = 60%

(2) การประมาณค่าแบบช่วง (Confidence Interval of Proportion)

สูตร

$$CI = \hat{p} \pm Z_{\alpha/2} \cdot \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$$

ตัวอย่าง จากตัวอย่างเดิม ($\hat{p} = 0.60$, $n = 200$, ระดับความเชื่อมั่น 95%, $Z = 1.96$)

$$\begin{aligned} CI &= 0.60 \pm 1.96 \cdot \sqrt{\frac{0.60(1 - 0.60)}{200}} \\ CI &= 0.60 \pm 1.96 \cdot \sqrt{\frac{0.24}{200}} = 0.60 \pm 1.96 \cdot \sqrt{0.0012} \\ CI &= 0.60 \pm 1.96 \cdot 0.0346 = 0.60 \pm 0.0678 \end{aligned}$$

ดังนั้น ช่วงความเชื่อมั่น 95% ของสัดส่วน = (0.532, 0.668) หรือ 53.2%–66.8%
การตีความ เรามั่นใจ 95% ว่าสัดส่วนจริงของนักเรียนที่ใช้สมาร์ทโฟนเพื่อการเรียนอยู่ระหว่าง 53.2%–66.8%

สรุปได้ว่า การประมาณค่าเป็นหัวใจของการวิเคราะห์เชิงอนุมาน เนื่องจากช่วยให้นักวิจัยสามารถจัดการกับความไม่แน่นอนของข้อมูลได้อย่างเป็นระบบ ทั้งการประมาณค่าเฉลี่ยและสัดส่วนในรูปแบบจุดและช่วง ต่างมีบทบาทในการให้ข้อมูลที่แม่นยำขึ้นเพื่อสนับสนุนการตัดสินใจเชิงวิชาการและการประยุกต์ใช้ในสังคมศาสตร์ได้อย่างมีประสิทธิภาพ

3. ช่วงความเชื่อมั่น (Confidence Interval)

ในการวิจัยเชิงปริมาณ นักวิจัยมักสนใจค่าพารามิเตอร์ของประชากร เช่น ค่าเฉลี่ย (μ) หรือสัดส่วน (p) แต่ไม่สามารถทราบค่าที่แท้จริงได้ เนื่องจากเก็บข้อมูลจากประชากรทั้งหมดเป็นไปไม่ได้ จึงต้องใช้ตัวอย่าง (sample) และคำนวณค่าประมาณ อย่างไรก็ตาม ค่าที่ได้จากตัวอย่างเป็นเพียงการ “ประมาณค่า” ซึ่งอาจคลาดเคลื่อนจากค่าจริงของประชากรได้ ดังนั้นสถิติอนุมานจึงพัฒนาแนวคิด ช่วงความเชื่อมั่น (Confidence Interval: CI) ขึ้นมา เพื่อช่วยให้นักวิจัยไม่เพียงได้ค่าประมาณ แต่ยังสามารถกำหนดช่วงของค่าที่เป็นไปได้ของพารามิเตอร์ พร้อมระบุระดับความเชื่อมั่น (confidence level) เช่น 90%, 95%, หรือ 99% (Frost, 2021; Penn State University, 2023)

ช่วงความเชื่อมั่น (Confidence Interval: CI) คือ ช่วงค่าที่คำนวณจากข้อมูลตัวอย่างซึ่งคาดว่าค่าพารามิเตอร์จริงของประชากรจะอยู่ในช่วงดังกล่าวด้วยความน่าจะเป็นตามที่กำหนด เช่น ช่วงความเชื่อมั่น 95% หมายความว่า หากทำการสุ่มตัวอย่างและสร้าง CI ซ้ำ 100 ครั้ง จะมีประมาณ 95 ครั้งในช่วงดังกล่าวครอบคลุมค่าพารามิเตอร์จริง (Cox, 2006)

สูตรการคำนวณ

1. ช่วงความเชื่อมั่นของค่าเฉลี่ย (Mean CI)

- กรณีทราบ σ

$$CI = \bar{X} \pm Z_{\alpha/2} \cdot \frac{\sigma}{\sqrt{n}}$$

- กรณีไม่ทราบ σ (ใช้ s และ t-distribution)

$$CI = \bar{X} \pm t_{\alpha/2, n-1} \cdot \frac{s}{\sqrt{n}}$$

2. ช่วงความเชื่อมั่นของสัดส่วน (Proportion CI)

$$CI = \hat{p} \pm Z_{\alpha/2} \cdot \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$$

ตัวอย่างการคำนวณ

ตัวอย่างที่ 1 ค่าเฉลี่ย (Mean CI)

นักวิจัยต้องการประมาณค่าเฉลี่ยรายได้ต่อวันของนักศึกษา

- ค่าเฉลี่ยตัวอย่าง ($\bar{X}$) = 500 บาท
- ส่วนเบี่ยงเบนมาตรฐานตัวอย่าง (s) = 50 บาท
- ขนาดตัวอย่าง (n) = 25

กรณีไม่ทราบ $\sigma \rightarrow$ ใช้ t-distribution (df = 24)

- ระดับความเชื่อมั่น 90% ($t_{0.025,24} \approx 1.711$)

$$CI = 500 \pm 1.711 \cdot \frac{50}{\sqrt{25}} = 500 \pm 1.711 \cdot 10 = 500 \pm 17.11$$

ช่วงความเชื่อมั่น 90% = (482.89, 517.11)

- ระดับความเชื่อมั่น 95% ($t_{0.025,24} \approx 2.064$)

$$CI = 500 \pm 2.064 \cdot 10 = 500 \pm 20.64$$

ช่วงความเชื่อมั่น 95% = (479.36, 520.64)

- ระดับความเชื่อมั่น 99% ($t_{0.025,24} \approx 2.797$)

$$CI = 500 \pm 2.797 \cdot 10 = 500 \pm 27.97$$

ช่วงความเชื่อมั่น 99% = (472.03, 527.97)

การตีความ

- หากต้องการความมั่นใจสูงขึ้น (เช่น 99%) ช่วงความเชื่อมั่นจะกว้างขึ้น
- หากยอมรับความเสี่ยงสูงขึ้น (เช่น 90%) ช่วงความเชื่อมั่นจะแคบลง

ตัวอย่างที่ 2 สัดส่วน (Proportion CI)

จากการสำรวจนักเรียน 200 คน พบว่า 120 คนใช้สมาร์ทโฟนเพื่อการเรียน

($\hat{p} = 0.60$)

ระดับความเชื่อมั่น 95% ($Z = 1.96$)

$$CI = 0.60 \pm 1.96 \cdot \sqrt{\frac{0.60(1 - 0.60)}{200}}$$

$$CI = 0.60 \pm 1.96 \cdot \sqrt{\frac{0.24}{200}} = 0.60 \pm 1.96 \cdot \sqrt{0.0012}$$

$$CI = 0.60 \pm 1.96 \cdot 0.0346 = 0.60 \pm 0.0678$$

ช่วงความเชื่อมั่น 95% = (0.532, 0.668) หรือ 53.2%–66.8%

การตีความ เรามั่นใจ 95% ว่าสัดส่วนจริงของนักเรียนที่ใช้สมาร์ทโฟนเพื่อการเรียนอยู่ระหว่าง 53.2%–66.8%

ข้อควรระวังในการตีความ CI

1. CI ไม่ได้หมายความว่ามีความน่าจะเป็น 95% ที่ค่าพารามิเตอร์จะอยู่ในช่วงนั้น แต่หมายถึง ถ้าสุ่มซ้ำหลายครั้ง จะมี 95% ของช่วงที่สร้างขึ้นครอบคลุมค่าพารามิเตอร์จริง

2. ความกว้างของ CI ขึ้นอยู่กับ

- ขนาดตัวอย่าง (n): n มาก $\rightarrow$ CI แคบลง
- ส่วนเบี่ยงเบนมาตรฐาน (σ หรือ s): การกระจายกว้าง $\rightarrow$ CI กว้างขึ้น

- ระดับความเชื่อมั่น: ยิ่งสูง → CI กว้างขึ้น

3. ควรเลือกความเชื่อมั่นตามบริบท เช่น งานสำรวจทั่วไปใช้ 95% แต่การวิจัยด้านการแพทย์ที่ความผิดพลาดมีผลรุนแรงอาจใช้ 99%

สรุปได้ว่า ช่วงความเชื่อมั่น (CI) เป็นเครื่องมือสำคัญที่ช่วยให้นักวิจัยเข้าใจความไม่แน่นอนของการประมาณค่า ทั้งค่าเฉลี่ยและสัดส่วน การเลือก CI ต้องพิจารณา ระดับความเชื่อมั่น ขนาดตัวอย่าง และความแปรปรวนของข้อมูล เพื่อให้ได้ข้อสรุปที่แม่นยำและเหมาะสมกับงานวิจัย

4. การทดสอบสมมติฐาน (Hypothesis Testing)

การทดสอบสมมติฐาน (Hypothesis Testing) เป็นหัวใจสำคัญของสถิติอนุมาน เพราะช่วยให้นักวิจัยสามารถใช้ข้อมูลจากตัวอย่างมาตัดสินใจว่า ข้อสมมติฐานเกี่ยวกับประชากร ควรถูกยอมรับหรือปฏิเสธ การทดสอบสมมติฐานทำให้การวิจัยทางสังคมศาสตร์มีความน่าเชื่อถือ เพราะไม่ได้อิงเพียงการสังเกต แต่เป็นการพิสูจน์ตามหลักสถิติ

4.1 หลักการทั่วไป

1) สมมติฐานศูนย์ (Null Hypothesis: H_0)

- ข้อสมมติฐานที่ตั้งขึ้นว่า “ไม่มีความแตกต่าง” หรือ “ไม่มีความสัมพันธ์” เช่น ค่าเฉลี่ยประชากรเท่ากับค่าหนึ่งที่กำหนด

- เขียนในรูป $H_0 : \mu = \mu_0$

2) สมมติฐานทางเลือก (Alternative Hypothesis: H_1)

- ข้อสมมติฐานที่ขัดแย้งกับ H_0

- เช่น $H_1 : \mu \neq \mu_0$ (ทดสอบสองด้าน) หรือ $H_1 : \mu > \mu_0$ (ทดสอบด้านเดียว)

3) ระดับนัยสำคัญ (Significance Level: α)

- ค่าความเสี่ยงในการปฏิเสธ H_0 ทั้งที่ H_0 เป็นจริง มักกำหนดที่ 0.05 หรือ 0.01

4) ค่าสถิติทดสอบ (Test Statistic)

- ค่าที่คำนวณจากข้อมูลตัวอย่างเพื่อตัดสินใจว่าจะปฏิเสธ H_0 หรือไม่

- เช่น Z-test หรือ t-test

5) ค่า p (p-value)

- ความน่าจะเป็นที่จะได้ผลการทดสอบอย่างที่เกิดขึ้นหรือมากกว่า ภายใต้เงื่อนไขว่า H_0 เป็นจริง (Frost, 2021)

4.2 ขั้นตอนการทดสอบสมมติฐาน

(Frost, 2021; Penn State University, 2023)

1) กำหนดสมมติฐาน H_0 และ H_1

- 2) กำหนดระดับนัยสำคัญ (α) เช่น 0.05
- 3) เลือกสถิติทดสอบ (Z-test หรือ t-test ขึ้นอยู่กับเงื่อนไข)
- 4) คำนวณค่าสถิติทดสอบ
- 5) หาค่า p-value หรือ Critical Value
- 6) ตัดสินใจ
 - ถ้า $p \leq \alpha \rightarrow$ ปฏิเสธ H_0
 - ถ้า $p > \alpha \rightarrow$ ไม่ปฏิเสธ H_0

ตัวอย่างการคำนวณ

ตัวอย่าง การทดสอบค่าเฉลี่ย (One-Sample t-test)

สมมติว่านักวิจัยต้องการทดสอบว่านักศึกษาที่มีคะแนนเฉลี่ยวิชาสถิติมากกว่า

70 หรือไม่

$$- H_0: \mu = 70$$

$$- H_1: \mu > 70$$

ข้อมูลจากตัวอย่าง

$$- \text{ค่าเฉลี่ยตัวอย่าง } (\bar{X}) = 73$$

$$- \text{ส่วนเบี่ยงเบนมาตรฐานตัวอย่าง } (s) = 8$$

$$- \text{ขนาดตัวอย่าง } (n) = 25$$

$$- \text{ระดับนัยสำคัญ } (\alpha) = 0.05$$

สูตร t-test

$$t = \frac{\bar{X} - \mu_0}{s/\sqrt{n}}$$

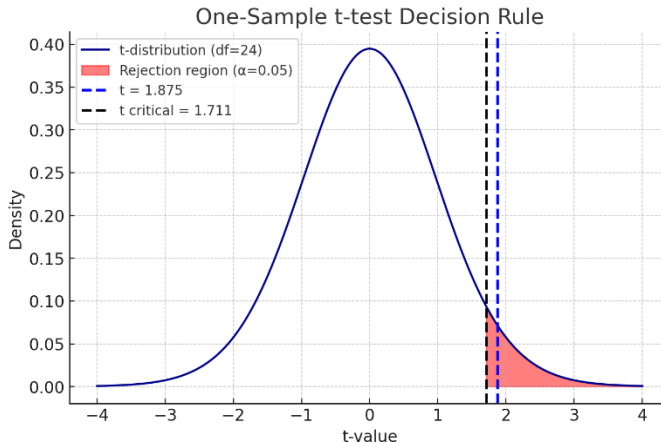
การแทนค่า

$$t = \frac{73 - 70}{8/\sqrt{25}} = \frac{3}{8/5} = \frac{3}{1.6} = 1.875$$

หาค่า Critical Value

$$- df = n - 1 = 24$$

$$- t_{0.05, 24} \approx 1.711 \text{ (ทดสอบด้านเดียว)}$$



แผนภาพที่ 5.1 การแจกแจง t-distribution

การตัดสินใจ

- $t = 1.875 > 1.711 \rightarrow$ ปฏิเสธ H_0

การตีความ

มีหลักฐานทางสถิติที่ระดับนัยสำคัญ 0.05 ว่าคะแนนเฉลี่ยของนักศึกษา มากกว่า 70 อย่างมีนัยสำคัญทางสถิติ

สรุปได้ว่า การทดสอบสมมติฐานเป็นกระบวนการที่ช่วยให้นักวิจัยใช้ข้อมูล ตัวอย่างในการตัดสินใจเกี่ยวกับประชากร โดยมีขั้นตอนชัดเจน ตั้งแต่การตั้งสมมติฐาน กำหนดระดับนัยสำคัญ เลือกสถิติทดสอบ คำนวณค่า test statistic และค่า p ไปจนถึง การตัดสินใจ การทดสอบสมมติฐานจึงเป็นเครื่องมือสำคัญที่ทำให้งานวิจัยเชิงปริมาณ สามารถยืนยันหรือหักล้างสมมติฐานได้อย่างเป็นระบบ

5. ค่าพี (p-value) และระดับนัยสำคัญ

เมื่อทำการทดสอบสมมติฐาน นักวิจัยจำเป็นต้องตัดสินใจว่าจะ “ปฏิเสธ” หรือ “ไม่ปฏิเสธ” สมมติฐานศูนย์ (H_0) ซึ่งการตัดสินใจนั้นมีความเสี่ยงที่จะผิดพลาด การใช้ค่า p (p-value) และระดับนัยสำคัญ (Significance Level: α) เป็นเครื่องมือช่วยให้การตัดสินใจมีหลักเกณฑ์ที่ชัดเจนและเป็นมาตรฐาน

1) ค่า p (p-value)

ค่า p คือ ความน่าจะเป็นที่จะได้ผลลัพธ์จากการทดสอบที่มีความสุดโต่งเท่ากับ หรือมากกว่าผลที่สังเกตได้จริง ภายใต้เงื่อนไขว่า H_0 เป็นจริง (Cox, 2006)

- ค่า p มีค่าอยู่ระหว่าง 0 ถึง 1

- ค่า p เล็ก หมายความว่า ผลลัพธ์ที่สังเกตได้ “ไม่น่าจะเกิดขึ้น” ถ้า H_0 เป็นจริง $\rightarrow$ มีเหตุผลเพียงพอที่จะปฏิเสธ H_0

- ค่า p ใหญ่ หมายความว่า ผลลัพธ์ที่สังเกตได้ “สอดคล้อง” กับ $H_0 \rightarrow$ ยังไม่มีหลักฐานเพียงพอที่จะปฏิเสธ H_0

ตัวอย่างเชิงความหมาย:

- $p = 0.03$ หมายความว่า หาก H_0 เป็นจริง โอกาสที่จะได้ผลลัพธ์สุดโต่งเช่นนี้หรือมากกว่าคือ 3%

- $p = 0.40$ หมายความว่า หาก H_0 เป็นจริง โอกาสที่จะได้ผลลัพธ์เช่นนี้หรือมากกว่ามีถึง 40% $\rightarrow$ จึงไม่ควรปฏิเสธ H_0

2) ระดับนัยสำคัญ (α)

ระดับนัยสำคัญ (Significance Level) คือ เกณฑ์ที่นักวิจัยกำหนดไว้ล่วงหน้าเพื่อใช้เปรียบเทียบกับค่า p และใช้ตัดสินใจในการปฏิเสธ H_0

- มักกำหนดที่ 0.05 (5%) หรือ 0.01 (1%)

- ค่า α = ความเสี่ยงที่ยอมรับได้ในการทำผิดพลาด แบบที่ 1 (Type I Error) หรือการปฏิเสธ H_0 ทั้งที่ H_0 เป็นจริง

ตัวอย่าง:

- ถ้า $\alpha = 0.05$ หมายความว่า นักวิจัยยอมรับความเสี่ยง 5% ที่จะตัดสินใจผิดพลาดในการปฏิเสธ H_0

ความสัมพันธ์ระหว่าง p และ α

- ถ้า $p \leq \alpha \rightarrow$ ปฏิเสธ $H_0 \rightarrow$ ผลการวิจัยมีนัยสำคัญทางสถิติ

- ถ้า $p > \alpha \rightarrow$ ไม่ปฏิเสธ $H_0 \rightarrow$ ผลการวิจัยไม่มีนัยสำคัญทางสถิติ

3) ตัวอย่างการคำนวณและการตีความ

ตัวอย่าง One-Sample t-test

นักวิจัยต้องการทดสอบว่านักศึกษาที่มีคะแนนเฉลี่ยวิชาสถิติแตกต่างจาก 70 หรือไม่

$$H_0: \mu = 70$$

$$H_1: \mu \neq 70$$

$$\text{ข้อมูล: } \bar{X} = 73, s = 8, n = 25, df = 24$$

สูตร

$$t = \frac{\bar{X} - \mu_0}{s/\sqrt{n}}$$

$$t = \frac{73 - 70}{8/\sqrt{25}} = \frac{3}{8/5} = \frac{3}{1.6} = 1.875$$

จากตาราง t (df = 24), $t = 1.875 \rightarrow p \approx 0.073$

- ถ้า $\alpha = 0.05 \rightarrow p (0.073) > 0.05 \rightarrow$ ไม่ปฏิเสธ H_0

- ถ้า $\alpha = 0.10 \rightarrow p (0.073) < 0.10 \rightarrow$ ปฏิเสธ H_0

การตีความ

- ที่ระดับนัยสำคัญ 0.05 $\rightarrow$ ยังไม่มีหลักฐานเพียงพอว่านักศึกษาที่มีคะแนนเฉลี่ยต่างจาก 70

- ที่ระดับนัยสำคัญ 0.10 $\rightarrow$ มีหลักฐานเพียงพอว่านักศึกษาที่มีคะแนนเฉลี่ยต่างจาก 70

การใช้ค่า p และ α ในงานวิจัยสังคมศาสตร์

- การสำรวจความคิดเห็นทางการเมือง: หากพบ $p < 0.05$ นักวิจัยสามารถสรุปได้ว่าความแตกต่างในระดับการสนับสนุนนโยบายระหว่างกลุ่มประชากรมีนัยสำคัญ

- การวิจัยทางการศึกษา: หากการทดลองสอนวิธีใหม่ให้ผลคะแนนสูงกว่าวิธีเดิมอย่างมีนัยสำคัญ ($p < 0.01$) สามารถยืนยันได้ว่าวิธีสอนใหม่นั้นมีประสิทธิภาพ

ข้อควรระวังในการตีความค่า p

1. ค่า p ไม่ได้บอกว่า H_0 เป็นจริงหรือไม่ แต่บอกว่าหาก H_0 เป็นจริง โอกาสที่จะได้ผลลัพธ์เช่นนี้มีมากน้อยเพียงใด

2. ค่า p ไม่ได้บอกขนาดของอิทธิพล (effect size) $\rightarrow$ ต้องวิเคราะห์เพิ่มเติมเพื่อดู “ความสำคัญเชิงปฏิบัติ” (practical significance) (Cohen, 1988)

3. การพึ่งพาเฉพาะค่า p อาจทำให้เกิดการตีความผิด $\rightarrow$ ควรใช้ร่วมกับช่วงความเชื่อมั่น (CI) และตัวชี้วัดอื่น ๆ

สรุปได้ว่า ค่า p (p-value) เป็นตัวบ่งบอกความน่าจะเป็นของผลลัพธ์ที่ได้ภายใต้สมมติฐานศูนย์ (H_0) ขณะที่ระดับนัยสำคัญ (α) เป็นเกณฑ์ที่นักวิจัยกำหนดล่วงหน้าเพื่อใช้ตัดสินใจ หากค่า $p \leq \alpha$ จะปฏิเสธ H_0 และสรุปว่าผลการวิจัยมีนัยสำคัญทางสถิติ แต่ถ้า $p > \alpha$ จะไม่ปฏิเสธ H_0 ทั้งนี้ การตีความผลควรพิจารณาพร้อมกับช่วงความเชื่อมั่น (CI) และขนาดอิทธิพล (effect size) เพื่อให้การวิจัยมีความหมายทั้งเชิงทฤษฎีและเชิงปฏิบัติ

6. การทดสอบค่าเฉลี่ย (Mean Testing)

การทดสอบค่าเฉลี่ยเป็นหนึ่งในกระบวนการทดสอบสมมติฐานที่ใช้บ่อยที่สุดในงานวิจัยทางสังคมศาสตร์ เศรษฐศาสตร์ และการบริหาร เนื่องจากค่าเฉลี่ย เป็นค่าที่บ่งบอกแนวโน้มเข้าสู่ศูนย์กลางของข้อมูล ตัวอย่างเช่น การตรวจสอบว่าคะแนนเฉลี่ยการเรียนของนักศึกษามหาวิทยาลัยแห่งหนึ่งสูงกว่าเกณฑ์มาตรฐานหรือไม่ การประเมินว่ารายได้เฉลี่ยของครัวเรือนในชนบทแตกต่างจากในเมืองหรือไม่ ฯลฯ

วิธีการทดสอบค่าเฉลี่ยหลัก ๆ มี 2 วิธี คือ Z-test และ t-test (One-Sample) โดยทั้งสองวิธีมีวัตถุประสงค์เหมือนกัน แตกต่างกันที่เงื่อนไขและข้อมูลที่มีอยู่

1. Z-test สำหรับค่าเฉลี่ย

Z-test ใช้เมื่อต้องการทดสอบค่าเฉลี่ยของประชากร โดยมีเงื่อนไขว่า **ทราบ** ส่วนเบี่ยงเบนมาตรฐานของประชากร (σ) หรือขนาดตัวอย่างใหญ่ ($n \geq 30$) ตามทฤษฎีบทขีดจำกัดกลาง (Central Limit Theorem) ซึ่งทำให้การแจกแจงของค่าเฉลี่ยตัวอย่างใกล้เคียงกับการแจกแจงปกติ (normal distribution)

สูตร

$$Z = \frac{\bar{X} - \mu_0}{\sigma/\sqrt{n}}$$

โดยที่

- $(\bar{X})$ = ค่าเฉลี่ยตัวอย่าง
- μ_0 = ค่าเฉลี่ยที่สมมติฐานศูนย์ (H_0) กำหนด
- σ = ส่วนเบี่ยงเบนมาตรฐานประชากร
- n = ขนาดตัวอย่าง

เงื่อนไขการใช้

- ใช้เมื่อทราบค่า σ หรือ n มีค่ามาก (≥ 30)
- ข้อมูลมีลักษณะใกล้เคียงการแจกแจงปกติ

ตัวอย่าง มีการอ้างว่า นักเรียนมัธยมมีเวลาอนเฉลี่ย 7 ชั่วโมง/วัน ($\mu_0 = 7$) นักวิจัยสุ่มตัวอย่างนักเรียน 36 คน พบว่า เวลาอนเฉลี่ย = 6.8 ชั่วโมง, $\sigma = 0.6$ ชั่วโมง, $\alpha = 0.05$

$$Z = \frac{6.8 - 7}{0.6/\sqrt{36}} = \frac{-0.2}{0.6/6} = \frac{-0.2}{0.1} = -2.0$$

- ค่า Z-critical (สองด้าน, $\alpha=0.05$) = ± 1.96
- $Z = -2.0 < -1.96 \rightarrow$ ปฏิเสธ H_0

การตีความ เวลารอนเฉลี่ยของนักเรียนแตกต่างจาก 7 ชั่วโมงอย่างมีนัยสำคัญทางสถิติที่ระดับ 0.05

2. t-test สำหรับค่าเฉลี่ย (One-Sample t-test)

t-test ใช้เมื่อต้องการทดสอบค่าเฉลี่ยของประชากร แต่ **ไม่ทราบค่า σ** และ **ขนาดตัวอย่างเล็ก ($n < 30$)** จึงต้องใช้การแจกแจง t ของ Student ซึ่งเหมาะสมกว่า Z-test เพราะ t-distribution มีหางหนากว่า ทำให้สะท้อนความไม่แน่นอนของการประมาณได้ดีกว่า

สูตร

$$t = \frac{\bar{X} - \mu_0}{s/\sqrt{n}}$$

โดยที่

- $(\bar{X})$ = ค่าเฉลี่ยตัวอย่าง
- μ_0 = ค่าเฉลี่ยที่กำหนดใน H_0
- S = ส่วนเบี่ยงเบนมาตรฐานตัวอย่าง
- n = ขนาดตัวอย่าง

เงื่อนไขการใช้

- ใช้เมื่อไม่ทราบค่า σ
- ขนาดตัวอย่างเล็ก ($n < 30$)
- ตัวอย่างต้องมาจากการแจกแจงที่ใกล้เคียงปกติ

ตัวอย่าง นักวิจัยต้องการทดสอบว่านักศึกษามีคะแนนเฉลี่ยวิชาสถิติต่างจาก 70 หรือไม่

- $H_0: \mu = 70$
- $H_1: \mu \neq 70$
- ข้อมูล: $\bar{X} = 73, s = 8, n = 25, df = 24, \alpha = 0.05$

$$t = \frac{73 - 70}{8/\sqrt{25}} = \frac{3}{8/5} = \frac{3}{1.6} = 1.875$$

- t-critical ($df=24$, สองด้าน, $\alpha=0.05$) $\approx \pm 2.064$
- $1.875 < 2.064 \rightarrow$ ไม่ปฏิเสธ H_0

การตีความ ไม่มีหลักฐานเพียงพอที่ระดับ 0.05 ว่าคะแนนเฉลี่ยของนักศึกษาแตกต่างจาก 70

ตารางที่ 5.1 การเลือกใช้ Z-test หรือ t-test

ประเด็น	Z-test	t-test
ความรู้เกี่ยวกับประชากร	ทราบค่า σ	ไม่ทราบค่า σ
ขนาดตัวอย่าง	$n \geq 30$	$n < 30$
การแจกแจง	ใช้การแจกแจงปกติ (Normal)	ใช้การแจกแจง t (Student's t)
ตัวอย่างการใช้	การสำรวจรายได้ประชากรจากสถิติแห่งชาติ	การเก็บข้อมูลจากนักศึกษาในกลุ่มเล็ก ๆ

3. ข้อควรระวังในการใช้การทดสอบค่าเฉลี่ย

1) การเลือกเครื่องมือผิดเงื่อนไข $\rightarrow$ ใช้ Z-test ทั้งที่ไม่ทราบค่า σ และตัวอย่างเล็ก อาจทำให้ผลคลาดเคลื่อน

2) การตีความผลลัพธ์ $\rightarrow$ ค่า p ที่เล็กบ่งชี้ว่าผลมีนัยสำคัญทางสถิติ แต่ไม่ได้หมายความว่า “นัยสำคัญทางปฏิบัติ (practical significance)”

3) สมมติฐานการแจกแจง (Normality Assumption) $\rightarrow$ ถ้าข้อมูลเบ้มากหรือตัวอย่างเล็ก การใช้ t-test อาจไม่เหมาะสม ต้องใช้สถิติแบบไม่อิงพารามิเตอร์ (non-parametric test)

สรุปได้ว่า Z-test และ t-test เป็นวิธีการพื้นฐานในการทดสอบค่าเฉลี่ย โดย Z-test เหมาะกับการกรณีที่ทราบค่าความเบี่ยงเบนมาตรฐานของประชากรหรือมีขนาดตัวอย่างใหญ่ ส่วน t-test เหมาะกับการที่ไม่ทราบค่า σ และมีขนาดตัวอย่างเล็ก อย่างไรก็ตาม ผลการทดสอบควรตีความร่วมกับค่า p-value ระดับนัยสำคัญ (α) และขนาดอิทธิพล (effect size) เพื่อให้ได้ข้อสรุปที่รอบด้านและน่าเชื่อถือ

7. ความผิดพลาดแบบที่ 1 และแบบที่ 2 (Type I & Type II Errors)

การทดสอบสมมติฐาน (Hypothesis Testing) มีวัตถุประสงค์เพื่อช่วยให้นักวิจัยตัดสินใจว่าข้อสมมติฐานที่ตั้งขึ้นเกี่ยวกับประชากร (Population) ควรถูก “ยอมรับ” หรือ “ปฏิเสธ” โดยอาศัยข้อมูลจากตัวอย่าง อย่างไรก็ตาม การตัดสินใจดังกล่าวอาจไม่ถูกต้องเสมอไป และอาจเกิด ความผิดพลาดทางสถิติ (Statistical Errors) ได้ ซึ่งแบ่งออกเป็นความผิดพลาดแบบที่ 1 (Type I Error) และความผิดพลาดแบบที่ 2 (Type II Error)

1. ความผิดพลาดแบบที่ 1 (Type I Error)

Type I Error คือ การปฏิเสธสมมติฐานศูนย์ (H_0) ทั้งที่ H_0 เป็นจริง เท่ากับการสรุปว่า “มีความแตกต่าง” หรือ “มีผล” ทั้งที่จริงแล้วไม่มี

สัญลักษณ์

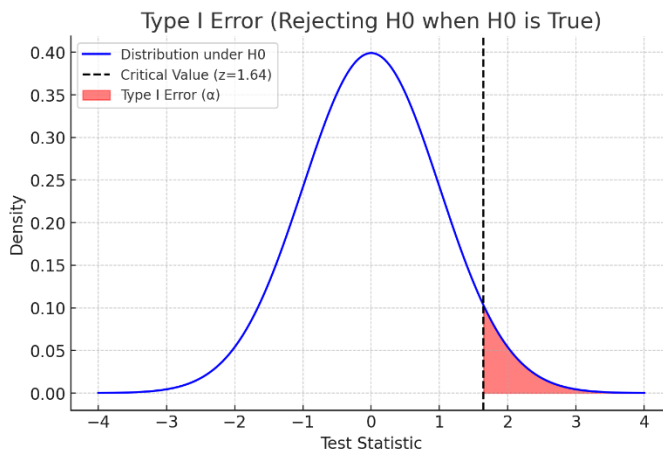
- โอกาสเกิดความผิดพลาดแบบที่ 1 = α (ระดับนัยสำคัญ)
- โดยทั่วไป นักวิจัยเลือก $\alpha = 0.05$ หรือ 0.01 ซึ่งหมายถึง ยอมรับความเสี่ยงในการตัดสินใจผิดพลาด 5% หรือ 1%

ตัวอย่าง งานวิจัยด้านการศึกษา สมมติว่า H_0 : คะแนนเฉลี่ยวิชาคณิตศาสตร์ของนักเรียน = 50

- นักวิจัยทดสอบแล้วพบค่า $p < 0.05 \rightarrow$ ปฏิเสธ H_0 และสรุปว่าคะแนนเฉลี่ยแตกต่างจาก 50

- แต่ในความเป็นจริง คะแนนเฉลี่ยของประชากร = 50 จริง $\rightarrow$ นักวิจัยเกิด Type I Error

อุปมา เหมือนกับการตัดสินใจว่าจำเลยมีความผิด ทั้งที่จริง ๆ แล้ว ไม่มีความผิด



แผนภาพที่ 5.2 Type I Error

2. ความผิดพลาดแบบที่ 2 (Type II Error)

Type II Error คือ การไม่ปฏิเสธ H_0 ทั้งที่ H_0 เป็นเท็จ เท่ากับการสรุปว่า “ไม่มีความแตกต่าง” หรือ “ไม่มีผล” ทั้งที่จริงแล้วมี

สัญลักษณ์

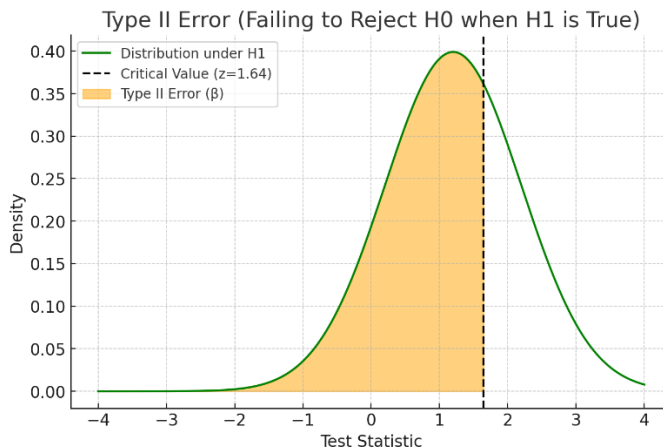
- โอกาสเกิดความผิดพลาดแบบที่ 2 = β
- ส่วนกลับของ β คือ Power of Test ($1-\beta$) ซึ่งหมายถึงความสามารถในการตรวจจับความแตกต่างได้อย่างถูกต้อง

ตัวอย่าง ในงานวิจัยเดียวกัน H_0 : คะแนนเฉลี่ยคณิตศาสตร์ = 50

- นักวิจัยคำนวณแล้วพบค่า $p > 0.05 \rightarrow$ ไม่ปฏิเสธ H_0 และสรุปว่าคะแนนเฉลี่ยไม่แตกต่างจาก 50

- แต่ความจริงคะแนนเฉลี่ยของประชากร = 55 ซึ่งแตกต่างจาก 50 อย่างมีนัยสำคัญ $\rightarrow$ นักวิจัยเกิด Type II Error

อุปมา เหมือนกับการตัดสินใจว่าจำเลยไม่มีความผิด ทั้งที่จริง ๆ แล้ว มีความผิด



แผนภาพที่ 5.3 Type II Error

3. วิธีลดความผิดพลาด

- ลด Type I Error (α)

เลือกใช้ระดับนัยสำคัญที่เข้มงวดขึ้น เช่น 0.01 แทน 0.05 แต่จะทำให้โอกาสเกิด Type II Error เพิ่มขึ้น

- ลด Type II Error (β)

1. เพิ่มขนาดตัวอย่าง (n) $\rightarrow$ ลดค่าความคลาดเคลื่อนมาตรฐาน (SE)
2. ใช้เครื่องมือวัดที่มีความเที่ยงตรง (validity) และเชื่อมั่นได้ (reliability)
3. กำหนดระดับนัยสำคัญ (α) ให้เหมาะสมกับบริบทการวิจัย
4. เลือกการทดสอบสถิติที่เหมาะสมกับลักษณะข้อมูล

4. ตัวอย่างจากงานวิจัยทางสังคมศาสตร์

- Type I Error: งานวิจัยสรุปว่านโยบายการศึกษาใหม่ทำให้ผลสัมฤทธิ์ของนักเรียนดีขึ้น แต่ในความเป็นจริงนโยบายดังกล่าวไม่มีผลใด ๆ

- Type II Error: งานวิจัยสรุปว่านโยบายการศึกษาใหม่ไม่ส่งผลต่อผลสัมฤทธิ์ของนักเรียน ทั้งที่จริงแล้วนโยบายนี้มีผลอย่างมีนัยสำคัญ

สรุปได้ว่า ความผิดพลาดแบบที่ 1 และแบบที่ 2 เป็นสิ่งที่เลี่ยงไม่ได้ในการทดสอบสมมติฐาน แต่สามารถจัดการได้ด้วยการเลือกระดับนัยสำคัญที่เหมาะสม การเพิ่มขนาดตัวอย่าง และการออกแบบการวิจัยอย่างรอบคอบ นักวิจัยจึงควรเข้าใจความหมายและผลกระทบของความผิดพลาดทั้งสองประเภท เพื่อหลีกเลี่ยงการสรุปผลที่ผิดพลาดและเพิ่มความน่าเชื่อถือของงานวิจัย

สรุปท้ายบท

การอนุมานทางสถิติ (Statistical Inference) เป็นกระบวนการสำคัญที่ทำให้ นักวิจัยสามารถนำข้อมูลจากกลุ่มตัวอย่างไปสรุปผลแทนประชากรได้อย่างมีประสิทธิภาพและน่าเชื่อถือ โดยอาศัยหลักความน่าจะเป็นเป็นพื้นฐานสำคัญ บทนี้ได้อธิบายตั้งแต่ ประชากร และตัวอย่าง ซึ่งเป็นจุดเริ่มต้นของการเก็บข้อมูล รวมถึงความแตกต่างระหว่าง พารามิเตอร์ และสถิติ ตลอดจนการเลือกวิธีสุ่มตัวอย่างทั้งแบบอิงความน่าจะเป็นและไม่อิงความน่าจะเป็น

จากนั้นกล่าวถึงการประมาณค่า ทั้งแบบจุดและแบบช่วง เพื่อใช้ข้อมูลจากตัวอย่างในการประมาณค่าพารามิเตอร์ของประชากร การสร้าง ช่วงความเชื่อมั่น ช่วยให้ นักวิจัยเข้าใจขอบเขตของค่าที่แท้จริงในระดับความเชื่อมั่นที่กำหนด หัวใจสำคัญอีก ประการหนึ่งคือ การทดสอบสมมติฐาน (Hypothesis Testing) ซึ่งประกอบด้วยขั้นตอน ชัดเจน ตั้งแต่การตั้งสมมติฐาน กำหนดระดับนัยสำคัญ เลือกสถิติทดสอบ คำนวณค่า test statistic และค่า p ไปจนถึงการตัดสินใจ โดยค่า p-value และระดับนัยสำคัญ (α) เป็น เกณฑ์สำคัญในการพิจารณาว่าจะปฏิเสธหรือไม่ปฏิเสธสมมติฐานศูนย์

นอกจากนี้ บทนี้ยังนำเสนอการทดสอบค่าเฉลี่ยผ่าน Z-test และ t-test (One-Sample) ซึ่งเป็นเครื่องมือพื้นฐานที่ใช้บ่อยในงานวิจัยทางสังคมศาสตร์ โดยมีเงื่อนไขในการเลือกใช้แตกต่างกัน พร้อมทั้งกล่าวถึง ความผิดพลาดแบบที่ 1 (Type I Error) และความผิดพลาดแบบที่ 2 (Type II Error) เพื่อให้ นักวิจัยตระหนักถึงข้อจำกัดและความเสี่ยงของการสรุปผล

กล่าวโดยสรุป การอนุมานทางสถิติเป็นทั้งศาสตร์และศิลป์ที่ช่วยให้นักวิจัยเชิงสังคมศาสตร์สามารถก้าวข้ามข้อจำกัดของข้อมูลตัวอย่างไปสู่การสร้างข้อสรุปที่มีนัยสำคัญต่อการทำความเข้าใจปรากฏการณ์ทางสังคม การนำวิธีการเหล่านี้ไปใช้อย่างถูกต้องและรอบคอบจะทำให้การวิจัยมีความน่าเชื่อถือและมีคุณค่าทั้งเชิงทฤษฎีและการประยุกต์ใช้

บทที่ 6

การวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร (ANALYSIS OF RELATIONSHIPS BETWEEN VARIABLES)

การวิจัยทางสังคมศาสตร์มีเป้าหมายสำคัญในการทำความเข้าใจความเชื่อมโยงระหว่างพฤติกรรมมนุษย์ เหตุการณ์ และโครงสร้างทางสังคม ซึ่งตัวแปรที่เกี่ยวข้องมักไม่ทำงานโดยลำพัง หากแต่มีความสัมพันธ์ระหว่างกันในระดับต่าง ๆ การวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร จึงเป็นกระบวนการสำคัญที่ช่วยให้นักวิจัยสามารถตรวจสอบและอธิบายรูปแบบการเปลี่ยนแปลงที่เกิดขึ้นได้อย่างเป็นระบบ

ตัวอย่างเช่น นักการศึกษาอาจสนใจตรวจสอบว่า จำนวนชั่วโมงการอ่านหนังสือของนักเรียนมีความสัมพันธ์กับคะแนนสอบหรือไม่ นักสังคมวิทยาอาจต้องการรู้ว่า รายได้เฉลี่ยของครัวเรือนสัมพันธ์กับระดับการศึกษาของสมาชิกในครอบครัวอย่างไร การศึกษาความสัมพันธ์ลักษณะนี้ไม่เพียงแต่ให้ข้อมูลเชิงพรรณนา แต่ยังสามารถนำไปใช้เชิงพยากรณ์ (prediction) ได้ เช่น การถดถอยเชิงเส้น (linear regression) ที่ใช้ในการทำนายค่าตัวแปรตามจากค่าตัวแปรอิสระ (Hair et al., 2010)

ดังนั้น การทำความเข้าใจเครื่องมือวิเคราะห์ความสัมพันธ์ ไม่ว่าจะเป็นค่าสหสัมพันธ์เพียร์สัน (Pearson's r) สำหรับข้อมูลเชิงปริมาณ ค่าสหสัมพันธ์สเปียร์แมน (Spearman's ρ) สำหรับข้อมูลเชิงอันดับ หรือการถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression) ล้วนเป็นพื้นฐานสำคัญของงานวิจัยเชิงปริมาณในสังคมศาสตร์ เพื่อสร้างข้อสรุปที่มีนัยสำคัญต่อการอธิบายและการกำหนดนโยบาย

1. ความหมายของความสัมพันธ์ (Correlation/ Relationship)

ในกระบวนการวิจัยทางสังคมศาสตร์ นักวิจัยมักสนใจตรวจสอบว่าตัวแปรต่าง ๆ มีความเชื่อมโยงกันอย่างไร เช่น รายได้มีผลต่อผลสัมฤทธิ์ทางการเรียนหรือไม่ ชั่วโมงการติวเสริมเกี่ยวข้องกับคะแนนสอบอย่างไร หรือความเชื่อมั่นทางการเมืองสัมพันธ์กับการเข้าร่วมกิจกรรมสาธารณะอย่างไร การตอบคำถามลักษณะนี้จำเป็นต้องใช้วิธีการทางสถิติที่เรียกว่า “การวิเคราะห์ความสัมพันธ์” (correlation analysis) ซึ่งช่วยให้เราทราบว่าตัวแปรหนึ่งเปลี่ยนแปลงไปพร้อมกับอีกตัวแปรหนึ่งหรือไม่ และถ้าใช่ มีความแน่นแฟ้นเพียงใด

1) ความหมายทางสถิติ

ความสัมพันธ์ (Correlation) ในทางสถิติ หมายถึง ระดับของการเปลี่ยนแปลงร่วมกันระหว่างตัวแปรตั้งแต่สองตัวขึ้นไป กล่าวคือ หากตัวแปรหนึ่งเปลี่ยนแปลง ตัวแปรอีกตัวก็มีแนวโน้มที่จะเปลี่ยนแปลงตามไปด้วยในทิศทางใดทิศทางหนึ่ง (สุภมาส อังศุโชติ, 2561) การวิเคราะห์ความสัมพันธ์จึงเป็นการตรวจสอบว่าตัวแปรต่าง ๆ เคลื่อนไหวไปด้วยกันหรือไม่ และถ้าใช่ มีความแน่นแฟ้นมากน้อยเพียงใด

Glass และ Hopkins (1984) อธิบายเพิ่มเติมว่า ความสัมพันธ์เป็นการพิจารณาการเปลี่ยนแปลงของตัวแปรสองตัวที่สัมพันธ์กันเชิงเส้นตรง โดยที่ค่าความสัมพันธ์ (correlation coefficient) เป็นตัวเลขที่บ่งบอกความเข้มข้นและทิศทางของความสัมพันธ์นั้น

2) ความสำคัญในการวิจัยสังคมศาสตร์

ในงานวิจัยด้านสังคมศาสตร์ การตรวจสอบความสัมพันธ์ระหว่างตัวแปรมีความสำคัญอย่างยิ่ง เพราะ

1) ช่วยอธิบายปรากฏการณ์ → เช่น รายได้ครอบคลุมความสัมพันธ์กับผลสัมฤทธิ์ทางการเรียน

2) ช่วยในการพยากรณ์ (prediction) → เช่น จำนวนชั่วโมงการติวเสริมสามารถทำนายคะแนนสอบได้

3) เป็นพื้นฐานของการวิเคราะห์ขั้นสูง → เช่น การถดถอยเชิงเส้นพหุคูณ (multiple regression), การวิเคราะห์เส้นทาง (path analysis)

สมถวิล วิจิตรวรรณ (2565) ชี้ว่า การเลือกใช้วิธีการวัดความสัมพันธ์ให้ถูกต้องตามระดับการวัดของข้อมูลเป็นสิ่งสำคัญมาก เพราะหากเลือกผิดอาจทำให้ผลที่ได้คลาดเคลื่อนไม่ตรงกับความเป็นจริง

3) ประเภทของความสัมพันธ์

(1) ความสัมพันธ์เชิงบวก (Positive Correlation): ตัวแปรหนึ่งเพิ่ม อีกตัวแปรก็มีแนวโน้มเพิ่ม เช่น ชั่วโมงอ่านหนังสือ (X) กับคะแนนสอบ (Y)

(2) ความสัมพันธ์เชิงลบ (Negative Correlation): ตัวแปรหนึ่งเพิ่ม อีกตัวแปรลดลง เช่น ระยะทางจากบ้านถึงโรงเรียน (X) กับความตรงต่อเวลา (Y)

(3) ไม่มีความสัมพันธ์ (No Correlation): ตัวแปรสองตัวไม่แสดงการเปลี่ยนแปลงร่วมกัน เช่น เบอร์รองเท้าของนักเรียน (X) กับคะแนนสอบ (Y)

4) ระดับความสัมพันธ์ (Strength of Correlation)

Best (1977) เสนอเกณฑ์สำหรับการแปลค่าความสัมพันธ์ (r) ที่ใช้แพร่หลาย

คือ

- 0.00 – 0.20 → ความสัมพันธ์ต่ำมาก (Very Low)
- 0.21 – 0.40 → ความสัมพันธ์ต่ำ (Low)
- 0.41 – 0.60 → ความสัมพันธ์ปานกลาง (Moderate)
- 0.61 – 0.80 → ความสัมพันธ์สูง (High)
- 0.81 – 1.00 → ความสัมพันธ์สูงมาก (Very High)

ตัวอย่างการตีความ

- $r = 0.18$ → ความสัมพันธ์ต่ำมาก สรุปรูปไม่ได้ว่ามีความเชื่อมโยง
- $r = 0.52$ → ความสัมพันธ์ปานกลาง สามารถใช้ประกอบการพยากรณ์ได้

ระดับหนึ่ง

- $r = 0.85$ → ความสัมพันธ์สูงมาก เชื่อว่าตัวแปรมีความใกล้ชิดกันอย่างมาก

5) ตัวอย่างการประยุกต์ใช้

1) การศึกษา: การวิจัยพบว่าจำนวนชั่วโมงอ่านหนังสือของนักเรียนมีความสัมพันธ์เชิงบวกสูงกับคะแนนสอบ O-NET ($r = 0.68$)

2) เศรษฐศาสตร์: รายได้ต่อครัวเรือนสัมพันธ์กับการใช้จ่ายเพื่อการศึกษา ($r = 0.72$)

3) รัฐศาสตร์: ความเชื่อมั่นในสถาบันการเมืองสัมพันธ์เชิงลบกับความถี่ในการเข้าร่วมการประท้วง ($r = -0.55$)

4) จิตวิทยา: ระดับความเครียดสัมพันธ์เชิงบวกกับการใช้สื่อสังคมออนไลน์ ($r = 0.47$)

6) ข้อควรระวังในการตีความ

Glass และ Hopkins (1984) เน้นว่า Correlation does not imply causation หรือความสัมพันธ์ไม่ใช่เหตุและผล ซึ่งเป็นข้อเข้าใจผิดที่พบบ่อย ตัวอย่างเช่น

- จำนวนร้านไอศกรีมกับจำนวนผู้จมน้ำมีค่าสหสัมพันธ์สูง เนื่องจากทั้งสองเพิ่มขึ้นพร้อมกันในทุกฤดูร้อน แต่ไม่ได้หมายความว่า การกินไอศกรีมทำให้จมน้ำ

- จำนวนเสาโทรศัพท์กับจำนวนผู้ป่วยมะเร็งอาจมีค่าสหสัมพันธ์สูง แต่มีตัวแปรแทรกซ้อน เช่น “ความหนาแน่นของประชากร”

ดังนั้น การตีความค่าความสัมพันธ์ต้องอาศัยทฤษฎีและบริบทเชิงสังคมมาประกอบ มิฉะนั้นอาจสรุปผิดพลาดได้

สรุปได้ว่า ความสัมพันธ์เป็นการตรวจสอบว่าตัวแปรสองตัวเปลี่ยนแปลงไปด้วยกันในทิศทางใดและมากน้อยเพียงใด ซึ่งเป็นพื้นฐานสำคัญของการวิเคราะห์เชิงอนุมานในสังคมศาสตร์ ช่วยอธิบายและทำนายปรากฏการณ์ได้ แต่ไม่สามารถสรุปเชิงเหตุผลได้โดยตรง การเลือกวิธีวิเคราะห์ความสัมพันธ์ต้องสอดคล้องกับระดับข้อมูลและเงื่อนไขการใช้ เพื่อให้ผลการวิจัยมีความน่าเชื่อถือ

2. สหสัมพันธ์เพียร์สัน (Pearson's Product-Moment Correlation, r)

ค่าสหสัมพันธ์เพียร์สัน (Pearson's r) เป็นสถิติที่ใช้วัด **ความสัมพันธ์เชิงเส้นตรง (linear relationship)** ระหว่างตัวแปรเชิงปริมาณ (interval/ratio scale) 2 ตัว ว่ามีการเปลี่ยนแปลงไปในทิศทางเดียวกันหรือตรงข้ามกัน และมีความแน่นแฟ้นมากน้อยเพียงใด

- ถ้า $r > 0 \rightarrow$ ความสัมพันธ์เชิงบวก (Positive) : X เพิ่ม $\rightarrow$ Y เพิ่ม
- ถ้า $r < 0 \rightarrow$ ความสัมพันธ์เชิงลบ (Negative) : X เพิ่ม $\rightarrow$ Y ลด
- ถ้า $r = 0 \rightarrow$ ไม่มีความสัมพันธ์เชิงเส้นตรง

ค่า r จะอยู่ระหว่าง **-1 ถึง +1** โดยค่าที่ใกล้ ± 1 หมายถึงความสัมพันธ์เข้มข้นมาก ส่วนค่าที่ใกล้ 0 หมายถึงความสัมพันธ์อ่อนหรือไม่มีเลย (Glass & Hopkins, 1984)

สูตรการคำนวณ

$$r = \frac{\sum(X_i - \bar{X}) - (Y_i - \bar{Y})}{\sqrt{\sum(X_i - \bar{X})^2 \cdot \sum(Y_i - \bar{Y})^2}}$$

หรือเขียนในรูปสมการผลรวมได้ว่า

$$r = \frac{n \sum XY - (\sum X)(\sum Y)}{\sqrt{[n \sum X^2 - (\sum X)^2] \cdot [n \sum Y^2 - (\sum Y)^2]}}$$

โดยที่

- X_i, Y_i = ค่าของตัวแปร X และ Y
- $\bar{X}, \bar{Y}$ = ค่าเฉลี่ยของ X และ Y
- n = จำนวนหน่วยตัวอย่าง

เงื่อนไขการใช้ Pearson's r

- 1) ตัวแปรเป็นเชิงปริมาณ (interval หรือ ratio scale)
- 2) ความสัมพันธ์ต้องมีลักษณะเชิงเส้นตรง (linear relationship)

3) ไม่มี outliers ที่รุนแรง เพราะ outlier เพียงตัวเดียวอาจทำให้ค่า r เปลี่ยน

4) การแจกแจงของข้อมูลใกล้เคียงปกติ (normal distribution) (สมถวิล วิจิตรวรรณ, 2565)

ตัวอย่างการคำนวณทีละขั้นตอน

ข้อมูล: ชั่วโมงอ่านหนังสือ (X) และคะแนนสอบ (Y) ของนักเรียน 5 คน

ตารางที่ 6.1 ชั่วโมงอ่านหนังสือและคะแนนสอบของนักเรียน

นักเรียน	X (ชั่วโมง)	Y (คะแนนสอบ)
1	2	65
2	3	70
3	4	75
4	5	85
5	6	90

ขั้นตอนที่ 1 หาค่าผลรวมที่ต้องใช้

$$- \sum X = 2+3+4+5+6 = 20$$

$$- \sum Y = 65+70+75+85+90 = 385$$

$$- \sum XY = (2 \times 65) + (3 \times 70) + (4 \times 75) + (5 \times 85) + (6 \times 90) = 1,605$$

$$- \sum X^2 = 2^2+3^2+4^2+5^2+6^2 = 90$$

$$- \sum Y^2 = 65^2+70^2+75^2+85^2+90^2 = 29,975$$

ขั้นตอนที่ 2 แทนค่าในสูตร

$$r = \frac{n \sum XY - (\sum X)(\sum Y)}{\sqrt{[n \sum X^2 - (\sum X)^2] \cdot [n \sum Y^2 - (\sum Y)^2]}}$$

$$r = \frac{5(1605) - (20)(385)}{\sqrt{[5(90) - (20)^2] \cdot [5(29,975) - (385)^2]}}$$

$$r = \frac{8,525 - 7,700}{\sqrt{(450) - (400)] \cdot (149,875) - (148,225)}}$$

$$r = \frac{325}{\sqrt{50 \times 1,650}} = \frac{325}{\sqrt{82,500}} = \frac{325}{287.25} \approx 0.99$$

ผลลัพธ์ $r \approx 0.99$

การตีความค่า r

ตามเกณฑ์ของ Best (1977)

$r = 0.99 \rightarrow$ ความสัมพันธ์เชิงบวกสูงมาก (very high positive correlation)

ความหมายในงานวิจัย ยิ่งนักเรียนอ่านหนังสือมาก ชั่วโมงการอ่านสัมพันธ์อย่างชัดเจนกับการได้คะแนนสอบสูงขึ้น

ข้อจำกัดของ Pearson's r

1) ใช้ได้เฉพาะความสัมพันธ์เชิงเส้นตรง หากความสัมพันธ์เป็นเส้นโค้ง (curvilinear) จะไม่สะท้อนความจริง

2) มีความอ่อนไหวต่อ outliers $\rightarrow$ ข้อมูลผิดปกติแม้เพียงค่าเดียวอาจทำให้ค่า r สูงหรือต่ำผิดไปจากความจริง

3) ไม่สามารถสรุปเชิงเหตุและผล (cause-effect) ได้ $\rightarrow$ แม้จะพบค่า r สูง ก็ไม่ได้หมายความว่า X เป็นต้นเหตุของ Y (Glass & Hopkins, 1984)

4) ไม่เหมาะกับข้อมูลแบบอันดับ (ordinal) หรือนามบัญญัติ (nominal) $\rightarrow$ ต้องใช้สถิติแบบไม่เป็นพารามेटริก เช่น Spearman's rho

สรุปได้ว่า Pearson's r เป็นสถิติพื้นฐานที่ใช้กันอย่างแพร่หลายในการวิจัยสังคมศาสตร์เพื่อวัดความสัมพันธ์เชิงเส้นตรงระหว่างตัวแปรสองตัวเชิงปริมาณ ค่าที่ได้ช่วยให้เห็นทั้งทิศทาง (บวก/ลบ) และความเข้มข้นของความสัมพันธ์ แต่การตีความต้องใช้ควบคู่กับทฤษฎีและบริบท เพราะความสัมพันธ์ไม่ใช่เหตุและผล การเลือกใช้ Pearson's r จึงต้องตรวจสอบเงื่อนไขของข้อมูลก่อนเสมอ เพื่อให้ผลการวิจัยมีความถูกต้องและน่าเชื่อถือ

3. สหสัมพันธ์แบบสเปียร์แมน (Spearman's Rank Correlation Coefficient, ρ)

สหสัมพันธ์แบบสเปียร์แมน (Spearman's rho หรือ ρ) เป็นสถิติ **ไม่เป็นพารามेटริก (non-parametric statistic)** ที่ใช้วัดความสัมพันธ์เชิงอันดับ (rank correlation) ระหว่างตัวแปรสองตัว เหมาะสำหรับกรณีที่

- ตัวแปรเป็นมาตราอันดับ (ordinal scale)
- ข้อมูลไม่เป็นปกติ (non-normal distribution)
- ความสัมพันธ์ไม่เป็นเชิงเส้นตรง (non-linear) แต่ยังคงเป็น *monotonic* (เมื่อ X เพิ่มขึ้น Y มีแนวโน้มเพิ่มหรือลดอย่างสม่ำเสมอ)

Spearman's rho จึงเป็นทางเลือกที่สำคัญเมื่อข้อมูลไม่สามารถใช้ Pearson's r ได้ (สมถวิล วิจิตรวรรณ, 2565; สุภมาส อังคุโชติ, 2561)

สูตรการคำนวณ

$$\rho = \frac{6\sum d_i^2}{n(n^2 - 1)}$$

โดยที่

- d_i = ผลต่างระหว่างอันดับของ X และ Y ของแต่ละคู่

- n = จำนวนคู่ข้อมูล

ถ้ามีค่าผูก (ties) จะต้องใช้วิธีการปรับอันดับเฉลี่ยก่อน

เงื่อนไขการใช้ Spearman's rho

1. ใช้กับข้อมูลแบบ ordinal หรือข้อมูลที่สามารถจัดอันดับได้
2. เหมาะกับกรณีที่ข้อมูลไม่เป็นปกติ (non-normal)
3. ใช้เมื่อข้อมูลมีความสัมพันธ์แบบ monotonic แม้ไม่เป็นเส้นตรงก็ได้
4. หากข้อมูลเป็นเชิงปริมาณแต่ไม่ผ่านสมมติฐานของ Pearson's r ก็มักใช้

Spearman แทน

ตัวอย่างการคำนวณ

ข้อมูล: จัดอันดับความถี่การใช้ห้องสมุด (X) และผลการประเมินความพึงพอใจ

(Y) ของนักศึกษา 6 คน

ตารางที่ 6.2 ความถี่การใช้ห้องสมุดและผลการประเมินความพึงพอใจของนักศึกษา

นักศึกษา	X (อันดับการใช้ห้องสมุด)	Y (อันดับความพึงพอใจ)	d = X-Y	d ²
1	1	2	-1	1
2	2	1	1	1
3	3	4	-1	1
4	4	3	1	1
5	5	5	0	0
6	6	6	0	0

$$-\sum d^2 = 1+1+1+1+0+0 = 4$$

$$-n = 6$$

แทนค่าในสูตร

$$\rho = 1 - \frac{6\sum d_i^2}{n(n^2 - 1)} = 1 - \frac{6(4)}{6(36 - 1)} = 1 - \frac{24}{210} = 1 - 0.114 \approx 0.886$$

$$\text{ผลลัพธ์: } \rho = 0.886$$

การตีความ

- ค่า $p = 0.886 \rightarrow$ ความสัมพันธ์เชิงบวกสูงมาก (very high positive correlation)

- แปลว่า นักศึกษาที่ใช้ห้องสมุดบ่อยมีแนวโน้มพึงพอใจสูงขึ้นด้วย

- ตามเกณฑ์การแปลค่าความสัมพันธ์ของ Best (1977) ค่านี้อยู่ในช่วง **สูงมาก**

ข้อจำกัดของ Spearman's rho

1. ใช้ได้เฉพาะข้อมูลที่สามารถจัดอันดับได้ หากข้อมูลเป็น nominal ไม่สามารถใช้ได้

2. แม้ใช้ได้กับความสัมพันธ์ monotonic ที่ไม่เป็นเส้นตรง แต่หากรูปแบบข้อมูลซับซ้อนมาก ค่า p อาจไม่สะท้อนความจริง

3. ไม่สามารถสรุปเชิงเหตุและผล เช่นเดียวกับ Pearson's r

4. หากมีค่าผูกอันดับ (ties) จำนวนมาก ความแม่นยำของค่า p อาจลดลง ต้องใช้วิธีปรับเฉลี่ยอันดับ

สรุปได้ว่า Spearman's rho เป็นสถิติที่เหมาะสมกับข้อมูลลำดับหรือข้อมูลที่ไม่เป็นปกติ ใช้เพื่อวัดความสัมพันธ์ในเชิง monotonic โดยไม่จำเป็นต้องเป็นเส้นตรง การคำนวณง่ายและตีความคล้าย Pearson's r โดยมีค่าอยู่ระหว่าง -1 ถึง +1 การเลือกใช้ Spearman's rho จึงเหมาะสำหรับงานวิจัยสังคมศาสตร์ที่ข้อมูลไม่เป็นพารามетริกหรือต้องการวิเคราะห์ข้อมูลเชิงอันดับ

4. การถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression)

การถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression) เป็นวิธีการทางสถิติที่ใช้ในการวิเคราะห์ **ความสัมพันธ์เชิงเส้นตรง** ระหว่างตัวแปรอิสระ (Independent Variable, X) 1 ตัว และตัวแปรตาม (Dependent Variable, Y) 1 ตัว จุดประสงค์สำคัญคือ

1. เพื่ออธิบายว่าตัวแปร X มีผลต่อการเปลี่ยนแปลงของตัวแปร Y อย่างไร

2. เพื่อสร้างสมการสำหรับการพยากรณ์ (prediction) ค่า Y จากค่า X

กล่าวโดยสรุป Regression เป็นการหาสมการเส้นตรงที่ “เหมาะสมที่สุด” (best fit line) ผ่านข้อมูล (Glass & Hopkins, 1984; Hair et al., 2010)

1) สมการถดถอย (Regression Equation)

สมการถดถอย (Regression Equation) คือ สมการเชิงเส้น ที่ใช้แสดงความสัมพันธ์ระหว่างตัวแปรอิสระ (Independent Variable, X) และตัวแปรตาม (Dependent Variable, Y) โดยแสดงในรูปแบบของเส้นตรง (straight line) ที่ “เหมาะสมที่สุด” (best fit line) ผ่านจุดข้อมูลทั้งหมด การสร้างสมการนี้มีเป้าหมายเพื่อ (1) อธิบายว่า

ค่าของ Y เปลี่ยนแปลงไปอย่างไรเมื่อค่า X เปลี่ยนแปลง และ (2) ใช้ในการทำนาย (prediction) ค่าของ Y จากค่า X

สมการทั่วไปคือ

$$Y = a + bX + \epsilon$$

โดยที่

- Y = ค่าตัวแปรตาม
- X = ค่าตัวแปรอิสระ
- a = ค่าคงที่ (intercept) $\rightarrow$ จุดที่เส้นตัดแกน Y เมื่อ $X = 0$
- b = สัมประสิทธิ์ถดถอย (regression coefficient) $\rightarrow$ แสดงการเปลี่ยนแปลงของ Y เมื่อ X เพิ่มขึ้น 1 หน่วย
- ϵ = ค่าความคลาดเคลื่อน (error term)

สูตรคำนวณสัมประสิทธิ์

1. สัมประสิทธิ์เชิงชัน (Slope, b)

$$b = \frac{n \sum XY - (\sum X)(\sum Y)}{n \sum X^2 - (\sum X)^2}$$

2. ค่าคงที่ (Intercept, a)

$$a = \frac{\sum Y - b \sum X}{n}$$

สูตรนี้ได้มาจากหลักการ **Least Squares Method** ซึ่งหาค่า a และ b ที่ทำให้ผลรวมกำลังสองของค่าคลาดเคลื่อนมีค่าน้อยที่สุด

สมการถดถอยคือหัวใจของการวิเคราะห์การถดถอยเชิงเส้นอย่างง่าย เพราะทำหน้าที่เชื่อมโยงระหว่างตัวแปรอิสระและตัวแปรตามในรูปแบบเส้นตรงที่เหมาะสมที่สุด การเข้าใจวิธีการคำนวณและการตีความค่าพารามิเตอร์ในสมการจึงเป็นทักษะสำคัญของนักวิจัยในสังคมศาสตร์

2) ค่าสัมประสิทธิ์ (Regression Coefficients)

ค่าสัมประสิทธิ์ถดถอย (Regression Coefficients) เป็นค่าที่ได้จากการประมาณสมการถดถอย โดยบ่งชี้ทิศทางและขนาดของผลกระทบ ที่ตัวแปรอิสระ (X) มีต่อตัวแปรตาม (Y)

ในสมการ

$$Y = a + bX + \epsilon$$

- **a (Intercept / จุดตัดแกน Y)**

ค่าคงที่ที่บอกว่า เมื่อ $X = 0$ ค่าเฉลี่ยของ Y จะเท่ากับ a

ใช้อธิบาย “ค่าพื้นฐาน” ของ Y ที่เกิดขึ้นโดยไม่ขึ้นกับ X

- **b (Slope / Regression Coefficient)**

ค่าที่บอกว่า เมื่อ X เพิ่มขึ้น 1 หน่วย ค่า Y จะเปลี่ยนไปเฉลี่ย b หน่วย

ใช้อธิบาย “อัตราการเปลี่ยนแปลง” ของ Y ตามการเปลี่ยนแปลงของ X

ประเภทของค่าสัมประสิทธิ์

1. ค่าคงที่ (Intercept, a)

- ถ้า a มีค่ามาก $\rightarrow$ แสดงว่าแม้ $X = 0$ ค่าเฉลี่ยของ Y ก็ยังคงสูง

- ถ้า a มีค่าน้อยหรือเป็นลบ $\rightarrow$ อาจตีความได้ว่าเมื่อ $X = 0$ ค่า Y จะต่ำ

หรือไม่สามารถเกิดขึ้นได้จริง (ต้องพิจารณาบริบท)

2. ค่าสัมประสิทธิ์เชิงชัน (Slope, b)

- ถ้า $b > 0 \rightarrow$ ความสัมพันธ์เชิงบวก (Positive relationship)

- ถ้า $b < 0 \rightarrow$ ความสัมพันธ์เชิงลบ (Negative relationship)

- ถ้า $b = 0 \rightarrow X$ ไม่มีผลต่อ Y

การตีความเชิงปฏิบัติ

- **เชิงปริมาณ b** บอกปริมาณการเปลี่ยนแปลงของ Y ต่อการเปลี่ยนแปลงของ X

- **เชิงคุณภาพ b** บอกทิศทางของความสัมพันธ์ (บวกหรือลบ)

- **เชิงพยากรณ์** สมการที่มี a และ b สามารถนำไปใช้ทำนายค่า Y เมื่อรู้ค่า X

ตัวอย่างการคำนวณและตีความ

ข้อมูล ชั่วโมงอ่านหนังสือ (X) และคะแนนสอบ (Y) ของนักเรียน 5 คน

ตารางที่ 6.3 ชั่วโมงอ่านหนังสือและคะแนนสอบของนักเรียน

นักเรียน	X (ชั่วโมง)	Y (คะแนนสอบ)
1	2	65
2	3	70
3	4	75
4	5	85
5	6	90

ขั้นตอนที่ 1 คำนวณค่าต่าง ๆ

- $\sum X = 20$
- $\sum Y = 385$
- $\sum XY = 1,605$
- $\sum X^2 = 90$
- $n = 5$

ขั้นตอนที่ 2 หาค่า ***b***

$$b = \frac{n \sum XY - (\sum X)(\sum Y)}{n \sum X^2 - (\sum X)^2} = \frac{5(1605) - (20)(385)}{5(90) - (20)^2}$$

$$b = \frac{8025 - 7700}{450 - 400} = \frac{325}{50} = 6.5$$

ขั้นตอนที่ 3 หาค่า ***a***

$$a = \frac{\sum Y - b \sum X}{n} = \frac{385 - 6.5(20)}{5} = \frac{385 - 130}{5} = \frac{255}{5} = 51$$

สมการถดถอยที่ได้

$$Y = 51 + 6.5X$$

การตีความผลลัพธ์

- ค่า ***a* = 51** → ถ้าไม่มีการอ่านหนังสือเลย ($X=0$) นักเรียนจะได้คะแนนสอบเฉลี่ย 51 คะแนน (เป็นค่าประมาณเชิงสถิติ)
- ค่า ***b* = 6.5** → เมื่อชั่วโมงการอ่านหนังสือเพิ่มขึ้น 1 ชั่วโมง คะแนนสอบเฉลี่ยจะเพิ่มขึ้น 6.5 คะแนน

การทำนาย (Prediction):

- ถ้านักเรียนอ่าน 4 ชั่วโมง → $Y = 51 + 6.5(4) = 77$ คะแนน
- ถ้านักเรียนอ่าน 7 ชั่วโมง → $Y = 51 + 6.5(7) = 96.5$ คะแนน

เงื่อนไขและข้อสมมติ (Assumptions)

1. ความสัมพันธ์ระหว่าง X และ Y ต้องเป็นเชิงเส้นตรง (linear relationship)
2. ค่าคลาดเคลื่อน (errors) ต้องมีการแจกแจงใกล้เคียงปกติ (normality)
3. ค่าคลาดเคลื่อนมีความแปรปรวนคงที่ (homoscedasticity)

4. ตัวแปรอิสระไม่สัมพันธ์กันมากเกินไป (multicollinearity) → ใช้กับ multiple regression

ข้อจำกัด

1. การถดถอยเชิงเส้นไม่สามารถสรุปว่า X เป็นต้นเหตุของ Y ได้ → ต้องพิจารณาทฤษฎีและบริบทประกอบ

2. ถ้าโครงสร้างข้อมูลไม่เป็นเชิงเส้นตรง ค่าสมการที่ได้อาจไม่สะท้อนความจริง

3. Outliers มีผลกระทบมากต่อเส้นถดถอย → ต้องตรวจสอบก่อนวิเคราะห์

4. ใช้กับข้อมูลตัวแปรปริมาณเป็นหลัก ข้อมูลที่เป็นเชิงคุณภาพต้องแปลงเป็น dummy variables ก่อน

ค่าสัมประสิทธิ์ถดถอย (a และ b) เป็นหัวใจของการวิเคราะห์การถดถอย เพราะบอกถึงค่าพื้นฐานของ Y และอัตราการเปลี่ยนแปลงของ Y ต่อการเปลี่ยนแปลงของ X การตีความต้องอาศัยทั้งมิติทางคณิตศาสตร์และมิติทางบริบทสังคมศาสตร์ เพื่อให้ผลลัพธ์มีความหมายที่ถูกต้องและนำไปใช้ประโยชน์ได้จริง

3) การทำนาย (Prediction)

การทำนาย (Prediction) ในการถดถอยเชิงเส้นอย่างง่าย หมายถึง การใช้สมการถดถอยที่ได้จากข้อมูลตัวอย่างมา **ประมาณค่าตัวแปรตาม (Y)** เมื่อทราบค่าของตัวแปรอิสระ (X) สมการที่สร้างขึ้นจึงมีบทบาทเป็นเครื่องมือพยากรณ์ ซึ่งเป็นหนึ่งในคุณค่าที่สำคัญของการวิเคราะห์การถดถอย (Hair et al., 2010)

เมื่อได้สมการถดถอยในรูป

$$Y = a + bX$$

โดยที่

- a = ค่าคงที่
- b = สัมประสิทธิ์ถดถอย
- X = ค่าตัวแปรอิสระที่เราต้องการใช้ในการทำนาย
- Y = ค่าพยากรณ์ของตัวแปรตาม

ขั้นตอนการทำนาย

1. สร้างสมการถดถอยจากข้อมูลตัวอย่าง
2. นำค่าตัวแปรอิสระ (X) ที่ต้องการทำนายมาแทนในสมการ
3. คำนวณค่า Y ที่ได้ → เรียกว่า **ค่าทำนาย (Predicted Value)**
4. ตีความผลลัพธ์ตามบริบทของการวิจัย

ตัวอย่างการคำนวณจริง

ข้อมูล ชั่วโมงอ่านหนังสือ (X) และคะแนนสอบ (Y) ของนักเรียน 5 คน จากการคำนวณ (ในหัวข้อที่ 2) ได้สมการถดถอยคือ

$$Y = 51 + 6.5X$$

การทำนาย

- ถ้านักเรียนอ่านหนังสือ 4 ชั่วโมง

$$Y = 51 + 6.5(4) = 51 + 26 = 77$$

→ คาดว่าจะได้คะแนนสอบ 77 คะแนน

- ถ้านักเรียนอ่านหนังสือ 7 ชั่วโมง

$$Y = 51 + 6.5(7) = 51 + 45.5 = 96.5$$

→ คาดว่าจะได้คะแนนสอบ ประมาณ 97 คะแนน

- ถ้านักเรียนอ่านหนังสือ 0 ชั่วโมง

$$Y = 51 + 6.5(0) = 51 + 0 = 51$$

→ คาดว่าจะได้คะแนนสอบ 51 คะแนน

การตีความเชิงปฏิบัติ

1. ค่า $a = 51$ แสดงว่าแม้นักเรียนไม่อ่านหนังสือเลย ก็ยังอาจได้คะแนนพื้นฐานประมาณ 51 คะแนน (อาจสะท้อนความรู้จากการเรียนในชั้นเรียน)

2. ค่า $b = 6.5$ แสดงว่าการอ่านหนังสือเพิ่มขึ้น 1 ชั่วโมง จะช่วยให้คะแนนสอบเพิ่มขึ้นเฉลี่ย 6.5 คะแนน

3. สมการนี้สามารถใช้ **วางกลยุทธ์ทางการศึกษา** เช่น กำหนดชั่วโมงอ่านหนังสือขั้นต่ำเพื่อให้ได้คะแนนตามเป้าหมาย

ข้อควรระวัง

1. การทำนายใช้ได้ดีในช่วงค่าของ X ที่ใกล้เคียงกับข้อมูลที่เก็บมา (interpolation) แต่ถ้าใช้เกินช่วง (extrapolation) เช่น ทำนายคะแนนสอบเมื่ออ่านหนังสือ 20 ชั่วโมง อาจคลาดเคลื่อนสูงเพราะข้อมูลจริงไม่ได้ครอบคลุมช่วงนั้น

2. สมการถดถอยไม่ได้ยืนยันเชิงสาเหตุ (causation) → แม้จะพบว่าการอ่านหนังสือสัมพันธ์กับคะแนนสอบ ก็ไม่สามารถสรุปได้ว่า “การอ่านหนังสือเพียงอย่างเดียว” เป็นต้นเหตุของคะแนนสอบที่สูงขึ้น

3. คุณภาพของการทำนายขึ้นอยู่กับความเหมาะสมของสมการ หากข้อมูลมี outliers หรือไม่เป็นเชิงเส้นตรง ผลการทำนายอาจไม่แม่นยำ

การทำนายด้วยสมการถดถอยเป็นขั้นตอนที่ทำให้สถิติไม่เพียงอธิบายปรากฏการณ์ แต่ยังสามารถใช้ **พยากรณ์** ปรากฏการณ์ได้ในอนาคต อย่างไรก็ตาม การ

ตีความต้องคำนึงถึงขอบเขตของข้อมูลและข้อจำกัด เพื่อให้การใช้สมการถดถอยเกิดประโยชน์สูงสุดต่อการวิจัยและการประยุกต์เชิงปฏิบัติ

สรุปได้ว่า การถดถอยเชิงเส้นอย่างง่ายเป็นเครื่องมือสำคัญที่ช่วยให้นักวิจัยสามารถอธิบายและทำนายความสัมพันธ์ระหว่างตัวแปรอิสระและตัวแปรตาม สมการที่ได้ช่วยให้เห็นทั้ง “ค่าคงที่” และ “อัตราการเปลี่ยนแปลง” ของตัวแปรตามเมื่อค่าของตัวแปรอิสระเปลี่ยนไป แม้เป็นวิธีที่ทรงพลัง แต่การตีความต้องอาศัยทฤษฎีและตรวจสอบสมมติฐานพื้นฐาน เพื่อให้ผลลัพธ์มีความน่าเชื่อถือและสะท้อนความจริงทางสังคมศาสตร์ได้อย่างถูกต้อง

5. การพล็อตกราฟแนวโน้ม (Trend Plotting)

การพล็อตกราฟแนวโน้ม หมายถึง การแสดงข้อมูลเชิงภาพเพื่อให้เห็นรูปแบบความสัมพันธ์ระหว่างตัวแปรอิสระ (X) และตัวแปรตาม (Y) โดยทั่วไปนิยมใช้ **กราฟกระจาย (Scatter Plot)** และวาด **เส้นแนวโน้ม (Trend Line)** หรือเส้นถดถอย (Regression Line) ลงไปในกราฟ เพื่อช่วยให้มองเห็นทิศทางและความเข้มข้นของความสัมพันธ์ได้อย่างชัดเจน (Glass & Hopkins, 1984; Hair et al., 2010)

ประเภทของกราฟที่ใช้

1. Scatter Plot (กราฟกระจาย)

ใช้แสดงข้อมูลเป็นจุด (X, Y) แต่ละคู่ → บอกทิศทางและรูปแบบความสัมพันธ์

2. Regression Line (เส้นถดถอย)

เป็นเส้นตรงที่ลากผ่านจุดข้อมูลโดยใช้วิธี Least Squares → แสดงแนวโน้มของข้อมูล

3. Residual Plot (กราฟเศษเหลือ)

ใช้ตรวจสอบคุณภาพของสมการถดถอย ว่ามีความเป็นเชิงเส้นและความแปรปรวนคงที่หรือไม่

ขั้นตอนการสร้างกราฟแนวโน้ม

1. เก็บข้อมูลของ X และ Y
2. สร้าง **Scatter Plot** โดยใช้ X บนแกนนอน และ Y บนแกนตั้ง
3. ตรวจสอบรูปแบบการกระจาย
 - จุดเรียงตัวจากซ้ายล่างไปขวาบน → ความสัมพันธ์เชิงบวก
 - จุดเรียงตัวจากซ้ายบนไปขวาล่าง → ความสัมพันธ์เชิงลบ
 - จุดกระจายไร้รูปแบบ → ไม่มีความสัมพันธ์

4. เพิ่ม เส้นแนวโน้ม (Trend Line) โดยใช้สมการถดถอย $Y = a + bX$

5. พิจารณาค่า R^2 (Coefficient of Determination) เพื่อดูว่าตัวแปรอิสระอธิบายความแปรปรวนของ Y ได้มากน้อยเพียงใด

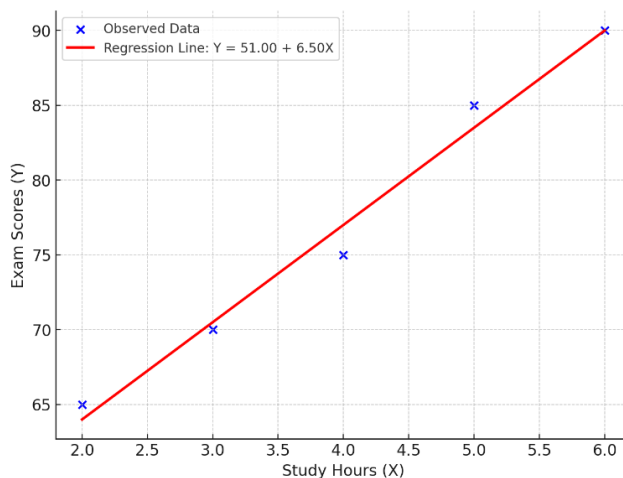
ตัวอย่างการพล็อตกราฟ

ข้อมูล ชั่วโมงอ่านหนังสือ (X) และคะแนนสอบ (Y) ของนักเรียน 5 คน

ตารางที่ 6.4 ตัวอย่างการพล็อตกราฟชั่วโมงอ่านหนังสือและคะแนนสอบของนักเรียน

นักเรียน	X (ชั่วโมง)	Y (คะแนนสอบ)
1	2	65
2	3	70
3	4	75
4	5	85
5	6	90

สมการถดถอย $Y = 51 + 6.5X$ (จากหน้า 135)



แผนภาพที่ 6.1 Scatter Plot + Regression Line

ภาพประกอบกราฟกระจาย (Scatter Plot) พร้อมเส้นถดถอยเชิงเส้น (Regression Line) แสดงความสัมพันธ์ระหว่าง ชั่วโมงอ่านหนังสือ (X) กับคะแนนสอบ (Y)

- จุดสีน้ำเงิน = ข้อมูลจริงแต่ละคู่
- เส้นสีแดง = สมการถดถอย $Y = 51 + 6.5X$

จากกราฟจะเห็นว่า จุดข้อมูลเรียงตัวใกล้เคียงเส้นตรง → แสดงความสัมพันธ์เชิงบวกที่ชัดเจน ระบุว่าชั่วโมงอ่านหนังสือที่มากขึ้นสัมพันธ์กับคะแนนสอบที่สูงขึ้น

การตีความ

- ค่าความสัมพันธ์ ($r \approx 0.99$) และค่า $R^2 \approx 0.98$ แสดงว่า ชั่วโมงอ่านหนังสือสามารถอธิบายคะแนนสอบได้ประมาณ 98%

- ยืนยันว่าการอ่านหนังสือมีผลเชิงบวกและชัดเจนต่อคะแนนสอบ

ประโยชน์ของการพล็อตกราฟแนวโน้ม

1. ช่วยตรวจสอบรูปแบบเบื้องต้น: มองเห็นทิศทางของข้อมูลชัดเจนก่อนทำการวิเคราะห์เชิงสถิติ

2. ค้นหาค่าผิดปกติ (Outliers) จุดข้อมูลที่เบี่ยงจากแนวโน้มจะมองเห็นได้ทันที

3. สื่อสารผลวิจัยได้ง่าย ผู้อ่านเข้าใจทิศทางและความสัมพันธ์ระหว่างตัวแปรโดยไม่ต้องตีความตัวเลขอย่างเดียว

4. สนับสนุนการตัดสินใจ กราฟแสดงให้เห็นว่าแนวโน้มที่ได้มีความสมเหตุสมผลหรือไม่

ข้อควรระวัง

1. กราฟความสัมพันธ์ไม่ใช่เหตุและผล (Correlation $\neq$ Causation) แม้เส้นแนวโน้มจะชี้ว่ามีความสัมพันธ์ แต่ไม่สามารถสรุปว่า X เป็นต้นเหตุของ Y ได้

2. Outliers มีอิทธิพลสูง จุดข้อมูลผิดปกติอาจบิดเบือนเส้นแนวโน้มและการตีความ

3. เลือกเส้นแนวโน้มให้เหมาะสม

- ถ้าข้อมูลเป็นเชิงเส้น $\rightarrow$ ใช้เส้นตรง (Linear Regression Line)

- ถ้าข้อมูลเป็นโค้ง $\rightarrow$ ต้องใช้ Polynomial Regression หรือ Nonlinear Regression

Regression

สรุปได้ว่า การพล็อตกราฟแนวโน้มเป็นวิธีการนำเสนอข้อมูลที่มีประสิทธิภาพช่วยให้เห็นความสัมพันธ์ระหว่างตัวแปรได้ง่ายขึ้น โดยเฉพาะการใช้ Scatter Plot ควบคู่กับเส้นถดถอย การตีความต้องพิจารณาทั้งค่า R^2 และทฤษฎีประกอบ เพื่อหลีกเลี่ยงการสรุปแบบผิดพลาดว่า “ความสัมพันธ์คือเหตุผล” ทั้งนี้ กราฟแนวโน้มจึงถือเป็นสะพานเชื่อมระหว่างข้อมูลเชิงตัวเลขกับความเข้าใจเชิงภาพ ในงานวิจัยสังคมศาสตร์

6. ข้อจำกัดและสมมติฐานของการวิเคราะห์

การวิเคราะห์ความสัมพันธ์และการถดถอยเป็นเครื่องมือที่สำคัญในการวิจัยสังคมศาสตร์ แต่การใช้เครื่องมือเหล่านี้มีทั้งข้อกำหนดเชิงสถิติ (assumptions) ที่ต้องตรวจสอบ และข้อจำกัด (limitations) ที่ต้องระวังในการตีความ หากละเลยอาจทำให้ผลวิเคราะห์คลาดเคลื่อนและข้อสรุปผิดพลาดได้

1) สมมติฐานของการวิเคราะห์

การวิเคราะห์ความสัมพันธ์และการถดถอยมีสมมติฐานทางสถิติ (Statistical Assumptions) ที่ต้องตรวจสอบก่อน หากสมมติฐานเหล่านี้ไม่เป็นจริง ผลลัพธ์อาจคลาดเคลื่อนและไม่สามารถตีความได้อย่างน่าเชื่อถือ

(1) สมมติฐานของสหสัมพันธ์เพียร์สัน (Pearson's r)

(1.1) ระดับการวัด (Scale of Measurement) ตัวแปรต้องอยู่ในระดับอันตรภาค (Interval) หรืออัตราส่วน (Ratio) เช่น คะแนนสอบ (Ratio scale) กับชั่วโมงอ่านหนังสือ (Interval scale) → ใช้ Pearson's r ได้

(1.2) ความเป็นเชิงเส้นตรง (Linearity) ความสัมพันธ์ระหว่างตัวแปร X และ Y ต้องอยู่ในรูปเส้นตรง หากเป็นโค้ง (Curvilinear) ค่า r จะไม่สะท้อนความจริง เช่น รายได้กับความสุข → Pearson's r ใช้ไม่ได้ ต้องใช้ Spearman's rho หรือ polynomial regression

(1.3) การแจกแจงใกล้เคียงปกติ (Normality) ตัวแปรควรมีการแจกแจงใกล้เคียงปกติ โดยเฉพาะในกรณีที่ใช้ทดสอบนัยสำคัญ ถ้าข้อมูลเบ้มาก ค่าความสัมพันธ์ที่ได้อาจไม่สะท้อนความจริง

(1.4) ไม่มีค่าผิดปกติรุนแรง (No extreme outliers) Outliers เพียงค่าเดียวสามารถทำให้ค่า r สูงหรือต่ำผิดปกติได้ เช่น เด็ก 1 คนที่อ่านหนังสือ 50 ชั่วโมงแต่สอบได้คะแนนต่ำ อาจทำให้ค่า r บิดเบือนไป (สมถวิล วิจิตรวรรณ, 2565)

(2) สมมติฐานของการถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression)

(2.1) Linearity ความสัมพันธ์ระหว่าง X และ Y ต้องเป็นเชิงเส้น สมการเส้นตรงต้องเป็นตัวแทนความสัมพันธ์ได้ เช่น ชั่วโมงอ่านหนังสือกับคะแนนสอบ

(2.2) Independence of Errors เศษเหลือ (residuals) ต้องเป็นอิสระจากกัน ใช้ Durbin-Watson test เพื่อตรวจสอบ เช่น ในข้อมูลแบบ time series

(2.3) Homoscedasticity ความแปรปรวนของ residuals ต้องเท่ากันในทุกค่าของ X หากไม่เท่ากันเรียกว่า heteroscedasticity ซึ่งทำให้ค่าประมาณไม่เสถียร

(2.4) Normality of Errors เศษเหลือต้องมีการแจกแจงใกล้เคียงปกติ (ตรวจสอบด้วย histogram หรือ Q-Q plot)

(2.5) No Multicollinearity (กรณี regression หลายตัว เช่น การศึกษา รายได้ และอาชีพสัมพันธ์กันมากเกินไป จะทำให้ค่าสัมประสิทธิ์ไม่เสถียร) ตัวแปรอิสระต้อง ไม่สัมพันธ์กันมากเกินไป มิฉะนั้นค่าพารามิเตอร์จะไม่น่าเชื่อถือ (Hair et al., 2010)

2) ข้อจำกัดของการวิเคราะห์

แม้การวิเคราะห์ความสัมพันธ์และการถดถอยจะเป็นเครื่องมือที่ทรงพลัง แต่ก็ มีข้อจำกัดที่นักวิจัยต้องตระหนัก ดังนี้

(1) Correlation $\neq$ Causation

ค่าสหสัมพันธ์สะท้อนเพียง “การเปลี่ยนแปลงร่วมกัน” ของตัวแปร แต่ไม่สามารถบอกได้ว่า X เป็นต้นเหตุของ Y หรือไม่ ตัวอย่างเช่น งานวิจัยด้านการศึกษาอาจ พบว่ารายได้ครอบครัวสัมพันธ์กับผลสัมฤทธิ์ทางการเรียน (Sirin, 2005) แต่ไม่อาจสรุปว่า รายได้เป็นสาเหตุโดยตรง เพราะมีตัวแปรแทรกซ้อน เช่น คุณภาพโรงเรียนหรือการ สนับสนุนจากครอบครัว

(2) ผลกระทบจาก Outliers

ข้อมูลผิดปกติแม้เพียงค่าเดียวอาจทำให้ค่า r หรือเส้นถดถอยเบี่ยงเบนไปไกล จากความจริง ตัวอย่างเช่น นักเรียนส่วนใหญ่มีคะแนนสอบระหว่าง 50-80 แต่มีนักเรียน คนหนึ่งได้ 100 คะแนน อาจทำให้ค่า r สูงเกินจริง

(3) ใช้ได้เฉพาะความสัมพันธ์เชิงเส้นตรง

หากความสัมพันธ์ระหว่างตัวแปรเป็นรูปโค้ง เช่น ความสุขกับรายได้ (ความสุข เพิ่มขึ้นตามรายได้ แต่ลดลงเมื่อรายได้เกินจุดหนึ่ง) Pearson's r และ regression แบบ เส้นตรงจะไม่สะท้อนความจริง ต้องใช้ polynomial regression หรือวิธีอื่นแทน

(4) การทำนายนอกช่วงข้อมูล (Extrapolation)

สมการถดถอยใช้ทำนายได้แม่นยำภายในช่วงของข้อมูล (interpolation) แต่ หากใช้เกินช่วง เช่น ใช้สมการพยากรณ์รายได้จากการศึกษาระดับ ป.ตรี เพื่อทำนายผู้ที่จบ ระดับ ป.เอก อาจได้ผลลัพธ์ที่คลาดเคลื่อน (Glass & Hopkins, 1984)

(5) ความซับซ้อนของปรากฏการณ์ทางสังคมศาสตร์

ตัวแปรทางสังคมมักมีปัจจัยแทรกซ้อนจำนวนมาก เช่น รายได้ การศึกษา สุขภาพ และทุนทางสังคม ต่างส่งผลต่อกันและกัน $\rightarrow$ การวิเคราะห์เพียงสองตัวแปรอาจ ไม่สะท้อนความจริง ต้องใช้วิธีวิเคราะห์ขั้นสูง เช่น Structural Equation Modeling (Norris, 2011)

(6) ความถูกต้องขึ้นอยู่กับคุณภาพข้อมูล

หากข้อมูลที่เก็บมาไม่เชื่อถือได้ (invalid) หรือไม่สอดคล้องกับทฤษฎี ค่าที่ได้ จากการวิเคราะห์ทางสถิติก็ไม่สามารถนำไปใช้ประโยชน์ได้จริง

3) ข้อเสนอแนะเชิงวิจัย

(1) ตรวจสอบสมมติฐานก่อนวิเคราะห์เสมอ เช่น ใช้ scatter plot เพื่อตรวจสอบความเป็นเชิงเส้น ตรวจสอบ residual plots เพื่อดู homoscedasticity

(2) ใช้สถิติที่เหมาะสมกับระดับข้อมูล เช่น ถ้าเป็นข้อมูลอันดับควรใช้ Spearman's rho แทน Pearson's r

(3) ตีความผลลัพธ์ร่วมกับทฤษฎีและบริบท ไม่ใช่ค่า r หรือสมการถดถอยอย่างเดียว

(4) หากต้องการสรุปเชิงสาเหตุ ควรใช้การวิจัยเชิงทดลอง (experimental research) หรือการวิเคราะห์ขั้นสูง เช่น path analysis หรือ SEM

สรุปได้ว่า การวิเคราะห์ความสัมพันธ์และการถดถอยเป็นเครื่องมือสำคัญในงานวิจัยทางสังคมศาสตร์ แต่มีข้อจำกัดหลายด้าน โดยเฉพาะเรื่อง **“Correlation $\neq$ Causation”** และผลกระทบจาก outliers การใช้วิธีการเหล่านี้อย่างถูกต้องต้องตรวจสอบสมมติฐานทางสถิติและพิจารณาบริบททางทฤษฎีควบคู่กัน การตระหนักถึงข้อจำกัดและสมมติฐานจะช่วยให้งานวิจัยมีความน่าเชื่อถือ และป้องกันการตีความผิดพลาดที่อาจส่งผลต่อข้อสรุปเชิงนโยบายหรือเชิงปฏิบัติ

สรุปท้ายบท

การวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรเป็นหัวใจสำคัญของงานวิจัยเชิงปริมาณในสังคมศาสตร์ เนื่องจากปรากฏการณ์ทางสังคมไม่ได้เกิดขึ้นโดยลำพัง แต่ตัวแปรต่าง ๆ มักเชื่อมโยงและส่งผลต่อกัน การทำความเข้าใจความสัมพันธ์จึงช่วยให้นักวิจัยสามารถอธิบายปรากฏการณ์ได้อย่างเป็นระบบ รวมทั้งใช้เป็นเครื่องมือในการพยากรณ์แนวโน้มที่อาจเกิดขึ้นในอนาคต

ภายในบทนี้ได้อธิบายตั้งแต่ความหมายและประเภทของความสัมพันธ์ ไม่ว่าจะเป็นความสัมพันธ์เชิงบวก เชิงลบ หรือไม่มีความสัมพันธ์ ตลอดจนการใช้ค่าสหสัมพันธ์เพียร์สัน (Pearson's r) สำหรับข้อมูลเชิงปริมาณที่มีความเป็นเชิงเส้นตรง และค่าสหสัมพันธ์สเปียร์แมน (Spearman's rho) สำหรับข้อมูลที่เป็นเชิงอันดับหรือข้อมูลที่ไม่เป็นไปตามสมมติฐานของเพียร์สัน ทั้งสองวิธีช่วยสะท้อนทิศทางและความแน่นแฟ้นของความสัมพันธ์ได้ แต่การเลือกใช้ต้องพิจารณาระดับข้อมูลและเงื่อนไขความเหมาะสมเสมอ

นอกจากนี้ บทนี้ยังได้อธิบายการถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression) ซึ่งเป็นการสร้างสมการเส้นตรงเพื่ออธิบายและทำนายค่าของตัวแปรตามจากตัวแปรอิสระ การตีความค่าสัมประสิทธิ์และค่าคงที่ในสมการถดถอยช่วยให้เข้าใจทั้งค่าพื้นฐานของตัวแปรตามและอัตราการเปลี่ยนแปลงที่เกิดจากตัวแปรอิสระ เครื่องมือนี้จึงเป็นประโยชน์

อย่างมากในเชิงการประยุกต์ใช้ โดยเฉพาะเมื่อผสานกับการพล็อตกราฟแนวโน้มน เช่น Scatter Plot และ Regression Line ที่ช่วยให้เห็นภาพความสัมพันธ์ได้ชัดเจนยิ่งขึ้น

อย่างไรก็ตาม การวิเคราะห์ความสัมพันธ์และการถดถอยมีทั้งสมมติฐานและข้อจำกัดที่นักวิจัยต้องตระหนัก สมมติฐานที่สำคัญ ได้แก่ ความเป็นเชิงเส้น การแจกแจงปกติ ความเป็นอิสระของเศษเหลือ และความแปรปรวนคงที่ ขณะเดียวกันก็มีข้อจำกัด เช่น ความสัมพันธ์ไม่ใช่เหตุผลโดยตรง (Correlation $\neq$ Causation) ผลกระทบจากข้อมูลผิดปกติ (outliers) การใช้ได้เฉพาะความสัมพันธ์เชิงเส้น และความเสี่ยงจากการทำนายนอกช่วงข้อมูล หากนักวิจัยละเลยปัจจัยเหล่านี้ อาจทำให้ข้อสรุปที่ได้ขาดความน่าเชื่อถือและไม่สามารถนำไปใช้เชิงปฏิบัติได้จริง

กล่าวโดยสรุป การวิเคราะห์ความสัมพันธ์และการถดถอยเป็นทั้งเครื่องมือในการอธิบายและพยากรณ์ปรากฏการณ์ทางสังคมศาสตร์ที่ทรงพลัง แต่การใช้ต้องควบคู่กับการตรวจสอบสมมติฐาน ความเข้าใจในข้อจำกัด และการตีความร่วมกับทฤษฎีและบริบททางสังคมศาสตร์ การตระหนักในจุดแข็งและข้อควรระวังเหล่านี้จะทำให้งานวิจัยมีความถูกต้อง น่าเชื่อถือ และสามารถนำไปประยุกต์ใช้เพื่อกำหนดนโยบายหรือแก้ไขปัญหาสังคมได้อย่างมีประสิทธิภาพ

บทที่ 7

การวิเคราะห์ความแตกต่างระหว่างกลุ่ม (ANALYSIS OF GROUP DIFFERENCES)

การวิจัยทางสังคมศาสตร์มักมีวัตถุประสงค์เพื่อทำความเข้าใจความแตกต่างที่เกิดขึ้นระหว่างกลุ่มของบุคคลหรือหน่วยข้อมูล เช่น การเปรียบเทียบผลสัมฤทธิ์ทางการเรียนระหว่างนักเรียนชายกับหญิง การศึกษาความแตกต่างของรายได้ระหว่างครอบครัวในเมืองกับชนบท หรือการวิเคราะห์ผลของการจัดการเรียนการสอนรูปแบบใหม่กับรูปแบบเดิม คำถามลักษณะนี้ไม่สามารถตอบได้ด้วยการอธิบายเชิงพรรณนาเพียงอย่างเดียว แต่จำเป็นต้องใช้ สถิติอนุมาน (Inferential Statistics) เพื่อวิเคราะห์ว่าความแตกต่างที่เห็นนั้น “มีนัยสำคัญทางสถิติ” หรือเกิดขึ้นเพียงเพราะความบังเอิญจากการสุ่มตัวอย่าง

ในทางสถิติ การเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่มเป็นเครื่องมือสำคัญในการตรวจสอบว่า ค่าเฉลี่ยของตัวแปรตาม (Dependent Variable) ของแต่ละกลุ่มมีความแตกต่างกันหรือไม่ และแตกต่างกันในระดับที่สามารถสรุปได้เชิงอนุมาน โดยมีเครื่องมือหลัก ได้แก่ t-test สำหรับกรณีมีกลุ่มเพียงสองกลุ่ม (อิสระหรือวัดซ้ำ) ANOVA (Analysis of Variance) สำหรับกรณีที่มีกลุ่มตั้งแต่สามกลุ่มขึ้นไป ซึ่งช่วยลดความเสี่ยงจากการใช้ t-test ซ้ำหลายครั้งที่อาจก่อให้เกิดความผิดพลาดแบบที่ 1 (Type I Error) และ Post-hoc tests เช่น Tukey HSD ซึ่งใช้ตรวจสอบเพิ่มเติมว่าคู่ของกลุ่มใดแตกต่างกัน หลังจากผล ANOVA พบว่ามีความแตกต่างโดยรวม

ดังนั้น บทนี้จะอธิบายแนวคิด วิธีการ และขั้นตอนการใช้สถิติในการเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม ตั้งแต่การวิเคราะห์เบื้องต้นด้วย t-test การใช้ ANOVA เพื่อตรวจสอบความแตกต่างหลายกลุ่ม ไปจนถึงการใช้ Post-hoc tests เพื่อระบุว่าคู่ของกลุ่มใดมีความแตกต่างที่ชัดเจน ตลอดจนการตรวจสอบสมมติฐานและข้อควรระวัง เพื่อให้ นักวิจัยสามารถเลือกใช้วิธีการที่ถูกต้องและตีความผลลัพธ์ได้อย่างน่าเชื่อถือ

1. การเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม

การวิจัยทางสังคมศาสตร์จำนวนมากตั้งคำถามว่า “กลุ่มต่าง ๆ มีผลลัพธ์แตกต่างกันหรือไม่” เช่น นักเรียนชายและหญิงมีผลสัมฤทธิ์ทางการเรียนต่างกันหรือไม่ ครอบครัวในเมืองกับชนบทมีรายได้ต่างกันหรือไม่ หรือผู้ที่เข้าร่วมโครงการอบรมมีทัศนคติทางการเมืองเปลี่ยนไปจากผู้ที่ไม่ได้เข้าร่วมหรือไม่ คำถามเหล่านี้ไม่สามารถตอบได้ด้วยการ

พรรณนาข้อมูลเพียงอย่างเดียว แต่ต้องใช้การวิเคราะห์เชิงสถิติที่เปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม เพื่อตอบคำถามเหล่านี้ นักวิจัยต้องเปรียบเทียบค่าเฉลี่ย (mean) ของตัวแปรตามในแต่ละกลุ่ม และใช้สถิติอนุมานเพื่อตัดสินว่าความแตกต่างที่พบ “เกิดจากความจริง” หรือ “เกิดจากความบังเอิญในการสุ่มตัวอย่าง” (Gravetter & Wallnau, 2009)

การเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม (Mean Comparison Between Groups) หมายถึง การใช้เครื่องมือทางสถิติในการตรวจสอบความแตกต่างของค่าเฉลี่ยของกลุ่มตัวอย่างตั้งแต่ 2 กลุ่มขึ้นไป โดยอาศัยหลักการแจกแจงความน่าจะเป็นและการทดสอบสมมติฐาน (Hypothesis Testing) (Howell, 2007)

- ถ้ามี 2 กลุ่ม → ใช้ *t-test*

- ถ้ามี 3 กลุ่มขึ้นไป → ใช้ *ANOVA* เพื่อหลีกเลี่ยงการใช้ *t-test* ซ้ำหลายครั้ง ซึ่งจะเพิ่มความเสี่ยงต่อ **Type I Error** (การปฏิเสธสมมติฐานศูนย์ H_0 ทั้งที่ H_0 เป็นจริง) (ปริญญา สิริอิตตะกุล, 2555)

การทดสอบเปรียบเทียบค่าเฉลี่ยตั้งอยู่บนสมมติฐาน 2 แบบคือ

- สมมติฐานศูนย์ (H_0): ค่าเฉลี่ยของทุกกลุ่มเท่ากัน → ไม่มีความแตกต่าง

- สมมติฐานทางเลือก (H_1): ค่าเฉลี่ยของกลุ่มอย่างน้อย 1 กลุ่มแตกต่าง

จากกลุ่มอื่น

ตัวอย่างเช่น

- $H_0: \mu_{ชาย} = \mu_{หญิง}$

- $H_1: \mu_{ชาย} \neq \mu_{หญิง}$

ตัวอย่างการคำนวณเบื้องต้น

นักวิจัยต้องการเปรียบเทียบรายได้เฉลี่ยต่อเดือนระหว่างคนในเมืองกับชนบท

ข้อมูล

- กลุ่มเมือง $\bar{X}_1 = 20,000$, $SD = 4,000$, $n = 40$

- กลุ่มชนบท $\bar{X}_2 = 15,000$, $SD = 3,000$, $n = 40$

ใช้ Independent *t-test*

$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{8_1^2}{n_1} + \frac{8_2^2}{n_2}}}$$

$$t = \frac{20,000 - 15,000}{\sqrt{\frac{(4,000)^2}{40} + \frac{(3,000)^2}{40}}} = \frac{5,000}{\sqrt{400,000 + 225,000}} = \frac{5,000}{\sqrt{625,000}}$$

$$t = \frac{5,000}{\sqrt{625,000}} = \frac{5,000}{790.57} = 6.33$$

ค่า df $\approx$ 78, t-critical (0.05) $\approx$ 1.99

$\rightarrow t = 6.33 > 1.99 \rightarrow$ ปฏิเสธ H_0

สรุป รายได้เฉลี่ยคนเมืองต่างจากชนบทอย่างมีนัยสำคัญ

ข้อควรระวัง

1. ความแตกต่างทางสถิติ $\neq$ ความแตกต่างทางปฏิบัติ บางครั้งผลต่างมีนัยสำคัญ แต่ขนาดผลต่าง (effect size) เล็กมากจนไม่มีความหมายเชิงนโยบาย

2. ต้องตรวจสอบสมมติฐานพื้นฐาน เช่น การแจกแจงปกติ (normality) และความเท่ากันของความแปรปรวน (homogeneity of variances) ด้วย Levene's Test

3. หลีกเลี่ยงการใช้ t-test หลายครั้ง เมื่อมีกลุ่มตั้งแต่ 3 ขึ้นไป ควรใช้ ANOVA

การประยุกต์ในงานวิจัยสังคมศาสตร์

- **ด้านเศรษฐศาสตร์** ศึกษาความแตกต่างของรายได้ระหว่างแรงงานภาคอุตสาหกรรม เกษตรกรรม และบริการ พบว่าค่าเฉลี่ยรายได้ต่างกันที่ระดับ .01

- **ด้านรัฐศาสตร์** เปรียบเทียบทัศนคติทางการเมืองของผู้ที่มีการศึกษา ม.ปลาย ปริญญาตรี และสูงกว่าปริญญาตรี พบว่ากลุ่มการศึกษาสูงกว่ามีระดับความเชื่อมั่นทางการเมืองมากกว่า

สรุปได้ว่า การเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่มเป็นพื้นฐานของการวิจัยเชิงปริมาณ ช่วยให้ตัดสินใจได้ว่าความแตกต่างที่สังเกตเกิดจากความจริงหรือเพียงความบังเอิญ การเลือกวิธีที่เหมาะสมขึ้นอยู่กับจำนวนกลุ่มและลักษณะข้อมูล หากวิเคราะห์อย่างถูกต้อง จะช่วยให้งานวิจัยมีความน่าเชื่อถือและสามารถนำไปใช้ในการกำหนดนโยบายหรือแนวทางการพัฒนาได้จริง

2. t-test สำหรับสองกลุ่มอิสระ (Independent Samples t-test)

Independent Samples t-test เป็นวิธีการทางสถิติที่ใช้เพื่อเปรียบเทียบค่าเฉลี่ยของตัวแปรตาม (dependent variable) ระหว่างสองกลุ่มที่เป็นอิสระต่อกัน เช่น กลุ่มนักเรียนชายและกลุ่มนักเรียนหญิง กลุ่มในเมืองและกลุ่มชนบท หรือกลุ่มทดลองและกลุ่มควบคุม

“อิสระต่อกัน” (Independent) หมายความว่า สมาชิกในแต่ละกลุ่มไม่ซ้ำกัน และการเลือกคนหนึ่งเข้ามาอยู่ในกลุ่มหนึ่ง จะไม่ส่งผลต่อโอกาสที่อีกคนจะอยู่ในอีกกลุ่มหนึ่ง

ความสำคัญของ Independent t-test คือการตอบคำถามวิจัยว่า “ค่าเฉลี่ยของสองกลุ่มมีความแตกต่างกันจริงหรือไม่” โดยไม่อาศัยเพียงการดูค่าเฉลี่ยและส่วน

เบี่ยงเบนมาตรฐาน แต่ใช้หลักการทดสอบสมมติฐานทางสถิติเพื่อลดความเสี่ยงในการสรุปผิดพลาด (Gravetter & Wallnau, 2009; Howell, 2007)

การทดสอบ t-test สำหรับสองกลุ่มอิสระใช้การตั้งสมมติฐานดังนี้

- สมมติฐานศูนย์ (Null Hypothesis: H_0): ค่าเฉลี่ยของทั้งสองกลุ่มไม่แตกต่างกัน

$$H_0: \mu_1 = \mu_2$$

- สมมติฐานทางเลือก (Alternative Hypothesis: H_1): ค่าเฉลี่ยของสองกลุ่มแตกต่างกัน

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

นอกจากนี้ยังอาจใช้การทดสอบด้านเดียว (one-tailed test) ถ้าสมมติว่าทิศทางชัดเจน เช่น

$$- H_1: \mu_1 > \mu_2$$

$$- H_1: \mu_1 < \mu_2$$

สมการการคำนวณ

1. กรณีความแปรปรวนเท่ากัน (Equal Variances)

$$t = \frac{\bar{x}_1 - \bar{x}_2}{s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

โดยที่

$$s_p = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}}$$

$\bar{x}_1, \bar{x}_2$ = ค่าเฉลี่ยของแต่ละกลุ่ม

s_1^2, s_2^2 = ความแปรปรวนของแต่ละกลุ่ม

n_1, n_2 = จำนวนตัวอย่างในแต่ละกลุ่ม

s_p = ส่วนเบี่ยงเบนมาตรฐานร่วม (pooled standard deviation)

2. กรณีความแปรปรวนไม่เท่ากัน (Unequal Variances, Welch's t-test)

$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

กรณีนี้ไม่ต้องใช้ s_p แต่ต้องประมาณค่า df ด้วยสูตรของ Welch (Field, 2013)

เงื่อนไขการใช้ (Assumptions)

1. ความเป็นอิสระของกลุ่ม → กลุ่มตัวอย่างทั้งสองต้องไม่ซ้ำกัน
2. มาตรวัดของตัวแปรตาม → ต้องเป็นมาตราอันตรภาค (interval scale) หรืออัตราส่วน (ratio scale) เช่น คะแนนสอบ, รายได้, น้ำหนัก
3. การแจกแจงใกล้เคียงปกติ (Normality) → แต่ละกลุ่มควรมีการแจกแจงปกติ โดยเฉพาะเมื่อ n มีค่าน้อยกว่า 30
4. ความแปรปรวนใกล้เคียงกัน (Homogeneity of variances) → ตรวจสอบด้วย Levene's Test
 - ถ้าผล Levene's Test มีค่า $p > .05$ → ถือว่าความแปรปรวนเท่ากัน ใช้สูตร Equal Variances
 - ถ้า $p < .05$ → ถือว่าความแปรปรวนไม่เท่ากัน ใช้สูตร Welch's t-test

ตัวอย่างการคำนวณ

นักวิจัยต้องการเปรียบเทียบคะแนนสอบวิชาสถิติของนักเรียนชายและหญิง

ข้อมูล

- ชาย ($n_1 = 30$): $\bar{x}_1 = 75$, $SD = 10$
- หญิง ($n_2 = 30$): $\bar{x}_2 = 80$, $SD = 12$

ขั้นตอนการคำนวณ (Welch's t-test)

1. คำนวณความแปรปรวน

- $s_1^2 = 10^2 = 100$
- $s_2^2 = 12^2 = 144$

2. คำนวณค่า t

$$t = \frac{75 - 80}{\sqrt{\frac{100}{30} + \frac{144}{30}}} = \frac{-5}{\sqrt{3.33 + 4.80}} = \frac{-5}{\sqrt{8.13}} = \frac{-5}{2.85} = -1.75$$

3. กำหนดค่า df และระดับนัยสำคัญ

ค่า $df \approx 58$, ที่ $\alpha = 0.05$ (two-tailed) → $t\text{-critical} \approx \pm 2.00$

4. การตัดสินใจ

$|t| = 1.75 < 2.00$ → ไม่ปฏิเสธ H_0

สรุป คะแนนเฉลี่ยของนักเรียนชายและหญิงไม่แตกต่างกันอย่างมีนัยสำคัญที่

ระดับ .05

ข้อควรระวัง

1. ต้องตรวจสอบ Normality และ Homogeneity of Variances ก่อนวิเคราะห์
2. ถ้าข้อมูลไม่เป็นปกติหรือมี outliers ควรใช้สถิติแบบไม่อิงพารามิเตอร์ เช่น

Mann–Whitney U Test

3. ต้องพิจารณา Effect Size (Cohen’s d) ควบคู่กับค่า p-value เพื่อบอก “ความสำคัญเชิงปฏิบัติ” (practical significance) ไม่ใช่แค่ความแตกต่างเชิงสถิติ

การประยุกต์ในงานวิจัยสังคมศาสตร์

- **ด้านเศรษฐศาสตร์** ศึกษาความแตกต่างของรายได้เฉลี่ยระหว่างแรงงานภาคอุตสาหกรรมและภาคเกษตร

- **ด้านรัฐศาสตร์** วิเคราะห์ทัศนคติทางการเมืองของผู้ที่เข้าร่วมอบรมพลเมืองเทียบกับผู้ที่ไม่ได้เข้าร่วม

สรุปได้ว่า Independent Samples t-test ใช้เปรียบเทียบค่าเฉลี่ยของสองกลุ่มที่อิสระต่อกัน เช่น ชาย-หญิง เมือง-ชนบท หรือกลุ่มทดลอง-ควบคุม เพื่อตรวจสอบว่าค่าเฉลี่ยต่างกันอย่างมีนัยสำคัญหรือไม่ ต้องอาศัยสมมติฐานเรื่องความเป็นอิสระ การแจกแจงปกติ และความแปรปรวนใกล้เคียงกัน หากไม่เป็นจริงควรใช้ Welch’s t-test หรือ Mann–Whitney U test ผลการวิเคราะห์ควรตีความร่วมกับค่า Effect Size เพื่อให้เข้าใจทั้งเชิงสถิติและเชิงปฏิบัติ

3. t-test สำหรับกลุ่มที่วัดซ้ำ (Paired Samples t-test)

t-test สำหรับกลุ่มที่วัดซ้ำ (Paired Samples t-test หรือ Dependent Samples t-test) ใช้เมื่อข้อมูลจากสองกลุ่ม ไม่เป็นอิสระต่อกัน (dependent groups) แต่มีความสัมพันธ์กันในเชิงโครงสร้างหรือวิธีการเก็บข้อมูล จุดมุ่งหมายคือการตรวจสอบว่าค่าเฉลี่ยของผลต่าง (difference scores) ระหว่างสองเงื่อนไขมีค่าแตกต่างจากศูนย์อย่างมีนัยสำคัญหรือไม่

ลักษณะการใช้ Paired t-test มีดังนี้

1. การวัดซ้ำ (Repeated Measures) → วัดตัวแปรเดียวกันจากกลุ่มเดียวกันในสองช่วงเวลา เช่น คะแนนสอบก่อนเรียน (pre-test) และหลังเรียน (post-test)

2. การจับคู่ (Matched Pairs) → จับคู่ตัวอย่างที่มีคุณสมบัติคล้ายกัน เช่น จับคู่เด็กเรียนตามเกรดเฉลี่ย แล้วสุ่มให้แต่ละคู่เข้ารับการสอนแบบต่างกัน

3. คู่ตามธรรมชาติ (Naturally Paired) → เช่น คู่สามีภรรยา พี่น้องหรือฝาแฝด

ความแตกต่างสำคัญระหว่าง Independent t-test กับ Paired t-test คือ Independent t-test เปรียบเทียบสองกลุ่มที่เป็นอิสระต่อกัน ขณะที่ Paired t-test ใช้กับข้อมูลที่มีการ “จับคู่” หรือ “ซ้ำวัด”

สมมติฐานการทดสอบ

- สมมติฐานศูนย์ (H_0): ค่าเฉลี่ยของผลต่างระหว่างการวัดทั้งสองครั้ง = 0

$$H_0: \mu_d = 0$$

- สมมติฐานทางเลือก (H_1): ค่าเฉลี่ยของผลต่าง $\neq 0$

$$H_1: \mu_d \neq 0$$

ทั้งนี้ หากต้องการทดสอบทิศทางเดียว สามารถกำหนด H_1 ว่า $\mu_d > 0$ หรือ $\mu_d < 0$ ได้ตามบริบทการวิจัย

สมการการคำนวณ

1. คำนวณผลต่างของแต่ละคู่

$$d_i = x_{1i} - x_{2i}$$

2. หาค่าเฉลี่ยของผลต่าง ($\bar{d}$) และส่วนเบี่ยงเบนมาตรฐานของผลต่าง (s_d)

$$\bar{d} = \frac{\sum d_i}{n}, s_d = \sqrt{\frac{\sum (d_i - \bar{d})^2}{n - 1}}$$

3. คำนวณค่า t

$$t = \frac{\bar{d}}{s_d / \sqrt{n}}$$

โดยที่

- $\bar{d}$ = ค่าเฉลี่ยผลต่าง
- s_d = ส่วนเบี่ยงเบนมาตรฐานของผลต่าง
- n = จำนวนคู่ข้อมูล
- $df = n - 1$

เงื่อนไขการใช้ (Assumptions)

1. การจับคู่/ซ้ำวัด ข้อมูลแต่ละคู่มีความสัมพันธ์กัน ไม่สามารถสุ่มแยกเป็นอิสระได้

2. ระดับการวัด ตัวแปรต้องเป็นอันตรภาค (interval) หรืออัตราส่วน (ratio)

3. การแจกแจงปกติของผลต่าง ผลต่างของคะแนน (difference scores) ควรมีการแจกแจงใกล้เคียงปกติ

4. ความถูกต้องของการวัด ใช้เครื่องมือเดียวกันในแต่ละรอบการเก็บข้อมูลเพื่อความน่าเชื่อถือ

ตัวอย่างการคำนวณทีละขั้นตอน

นักวิจัยต้องการทดสอบว่าวิธีสอนใหม่ช่วยเพิ่มคะแนนสอบสถิติของนักเรียนหรือไม่ โดยสุ่มนักเรียน 5 คนมาทำแบบทดสอบก่อนเรียนและหลังเรียน

ตารางที่ 7.1 คะแนนทดสอบก่อนเรียนและหลังเรียน

นักเรียน	Pre-test	Post-test	d = Pre – Post
1	60	70	-10
2	65	75	-10
3	70	78	-8
4	62	72	-10
5	68	74	-6

- คำนวณค่าเฉลี่ยผลต่าง

$$\bar{d} = \frac{-10 + -10 + -8 + -10 + -6}{5} = -8.8$$

- คำนวณส่วนเบี่ยงเบนมาตรฐานของผลต่าง

$$S_d \approx 1.79$$

- คำนวณค่า t

$$t = \frac{-8.8}{1.79/\sqrt{5}} = \frac{-8.8}{1.79} = -11.0$$

- df = 5 – 1 = 4

- t-critical ($\alpha = 0.05$, two-tailed, df = 4) $\approx \pm 2.776$

ผลการทดสอบ $|t| = 11.0 > 2.776 \rightarrow$ ปฏิเสธ H_0

สรุป คะแนนสอบหลังเรียนสูงกว่าก่อนเรียนอย่างมีนัยสำคัญทางสถิติที่ระดับ

.05

ข้อควรระวัง

1. อิทธิพลของ outliers ผลต่างที่ผิดปกติอาจทำให้ค่าเฉลี่ยผลต่างเพี้ยนและค่า t-test ไม่สะท้อนความจริง

2. การแจกแจงปกติ หากผลต่างไม่เป็นปกติ ควรใช้สถิติไม่อิงพารามิเตอร์ เช่น Wilcoxon Signed-Rank Test

3. Practical vs Statistical significance แม้ผลต่างมีนัยสำคัญทางสถิติ แต่ควรตรวจสอบด้วย Effect Size (Cohen's d) เพื่อบอกว่าสำคัญในเชิงปฏิบัติหรือไม่

สรุปได้ว่า Paired Samples t-test ใช้เปรียบเทียบค่าเฉลี่ยของกลุ่มที่มีความสัมพันธ์กัน เช่น การวัดซ้ำ (pre-post) หรือกลุ่มที่จับคู่กัน เพื่อตรวจสอบว่าค่าเฉลี่ยของผลต่างแตกต่างจากศูนย์อย่างมีนัยสำคัญหรือไม่ เหมาะกับงานวิจัยที่ต้องการวัดการเปลี่ยนแปลงภายในกลุ่มเดียวกัน อย่างไรก็ตาม ต้องตรวจสอบเงื่อนไข เช่น การแจกแจงปกติของผลต่าง และระวังผลกระทบจาก outliers ผลการทดสอบควรตีความร่วมกับค่า Effect Size (Cohen's d) เพื่อให้เห็นทั้งนัยสำคัญเชิงสถิติและเชิงปฏิบัติ

4. การวิเคราะห์ความแปรปรวน (Analysis of Variance: ANOVA)

1) การทดสอบ F (F-test)

การทดสอบ F เป็นสถิติที่ใช้ในการวิเคราะห์ความแปรปรวน (Analysis of Variance: ANOVA) เพื่อเปรียบเทียบค่าเฉลี่ยของกลุ่มตั้งแต่ 3 กลุ่มขึ้นไป ว่ามีความแตกต่างกันอย่างมีนัยสำคัญหรือไม่ (Howell, 2007) หลักการคือการเปรียบเทียบ ความแปรปรวนระหว่างกลุ่ม (between groups variance) กับความแปรปรวนภายในกลุ่ม (within groups variance) ถ้าไม่มีความแตกต่างจริง ความแปรปรวนทั้งสองควรมีค่าใกล้เคียงกัน แต่ถ้ามีความแตกต่าง ความแปรปรวนระหว่างกลุ่มจะสูงกว่าอย่างมีนัยสำคัญ (Gravetter & Wallnau, 2009) ดังนั้น การทดสอบ F คือ การหาสัดส่วนของความแปรปรวนระหว่างกลุ่มต่อความแปรปรวนภายในกลุ่ม

สมมติฐานในการทดสอบ F

- สมมติฐานศูนย์ (H_0): ค่าเฉลี่ยของทุกกลุ่มเท่ากัน

$$\mu_1 = \mu_2 = \mu_3 \dots = \mu_k$$

- สมมติฐานทางเลือก (H_1): ค่าเฉลี่ยของอย่างน้อยหนึ่งกลุ่มแตกต่างจากกลุ่มอื่น

$$\mu_i \neq \mu_j \text{ สำหรับบางคู่ของ } i \text{ และ } j$$

สมการการคำนวณค่า F

$$F = \frac{MS_{between}}{MS_{within}}$$

โดยที่

- MS = Mean Square (ค่าเฉลี่ยความแปรปรวน)

$$MS_{between} = \frac{SS_{between}}{df_{between}}$$

$$MS_{within} = \frac{SS_{within}}{df_{within}}$$

คำนิยาม

- SS (Sum of Squares) ผลรวมกำลังสองของผลต่าง
- df (Degrees of Freedom) องศาอิสระ
- Between groups ความแปรปรวนที่เกิดจากความต่างของค่าเฉลี่ยแต่ละกลุ่ม
- Within groups (Error) ความแปรปรวนที่เกิดจากความแตกต่างภายในกลุ่ม

ขั้นตอนการวิเคราะห์ F-test (ANOVA)

1. หาค่าเฉลี่ยของแต่ละกลุ่ม และค่าเฉลี่ยรวม (Grand Mean)
2. คำนวณ Sum of Squares
 - $SS_{between} = \sum n_i (\bar{x}_i - \bar{x})^2$
 - $SS_{within} = \sum (\bar{x}_{ij} - \bar{x})^2$
3. หาค่า df
 - $df_{between} = k - 1$
 - $df_{within} = N - k$
4. คำนวณ Mean Squares (MS)
 - $MS_{between} = SS_{between} / df_{between}$
 - $MS_{within} = SS_{within} / df_{within}$
5. คำนวณค่า F
 - $F = MS_{between} / MS_{within}$
6. หาค่า p-value จากตาราง F หรือซอฟต์แวร์สถิติ
7. ตัดสินใจ
 - ถ้า $p \leq 0.05 \rightarrow$ ปฏิเสธ $H_0 \rightarrow$ มีความแตกต่างระหว่างกลุ่ม
 - ถ้า $p > 0.05 \rightarrow$ ยอมรับ $H_0 \rightarrow$ ไม่พบความแตกต่าง

ตัวอย่างการคำนวณ F-test

นักวิจัยต้องการเปรียบเทียบผลสัมฤทธิ์ทางการเรียนของนักเรียน 3 กลุ่มที่ใช้

วิธีการสอนต่างกัน

ตารางที่ 7.2 ผลสัมฤทธิ์ทางการเรียนของนักเรียน

กลุ่ม	คะแนน
A	70, 75, 80
B	65, 68, 72
C	85, 90, 95

ขั้นตอนที่ 1 ค่าเฉลี่ย

- กลุ่ม A: $\bar{x}_A = \frac{70+75+80}{3} = 75$
- กลุ่ม B: $\bar{x}_B = \frac{65+68+72}{3} = 68.3$
- กลุ่ม C: $\bar{x}_C = \frac{85+90+95}{3} = 90$
- ค่าเฉลี่ยรวม (Grand Mean) = $\bar{x} = \frac{75+68.33+90}{3} = 77.8$

ขั้นตอนที่ 2 $SS_{within} = \sum(\bar{x}_{ij} - \bar{x})^2$

- กลุ่ม A: $(70-75)^2 + (75-75)^2 + (80-75)^2 = 50$
- กลุ่ม B: $(65-68.3)^2 + (68-68.3)^2 + (72-68.3)^2 = 24.6$
- กลุ่ม C: $(85-90)^2 + (90-90)^2 + (95-90)^2 = 50$
- $SS_{within} = 50 + 24.6 + 50 = 124.6$

ขั้นตอนที่ 3 $SS_{between} = \sum n_i(\bar{x}_i - \bar{x})^2$ โดย $n_i = 3$

- $SS_{between} = 3(75-77.8)^2 + 3(68.3-77.8)^2 + 3(90-77.8)^2 \approx 738.9$
- $SS_{total} = SS_{between} + SS_{within} = 738.9 + 124.7 \approx 863.4$

ขั้นตอนที่ 4 (df), MS และ F

- $df_{between} = 3 - 1 = 2$
- $df_{within} = 9 - 3 = 6$
- $MS_{between} = 738.9/2 = 369.4$
- $MS_{within} = 124.7/6 = 20.8$
- $F = 369.4/20.8 = 17.8$

ขั้นตอนที่ 5 การตัดสินใจ

- ที่ $df = (2, 6)$, $\alpha = 0.05 \rightarrow F_{critical} \approx 5.14$
- ค่า $F = 17.8 > 5.14 \rightarrow$ ปฏิเสธ H_0

สรุป คะแนนเฉลี่ยของนักเรียน 3 กลุ่มมีความแตกต่างกันอย่างมีนัยสำคัญ

ตารางที่ 7.3 ผลการเปรียบเทียบผลสัมฤทธิ์ทางการเรียน

Source	SS	df	MS	F	p-value	F_crit ($\alpha=0.05$)	Decision
Between Groups	738.9	2	369.4	17.8	0.0	5.1	Reject H_0
Within Groups	124.7	6	20.8				
Total	863.6	8					

หมายเหตุ: ที่ $df = (2,6)$ และระดับนัยสำคัญ 0.05 ค่า $F_{crit} = 5.1$ และค่า F ที่คำนวณได้ = 17.8 ซึ่งมากกว่าเกณฑ์วิกฤติ จึงปฏิเสธสมมติฐานศูนย์ และสรุปได้ว่าคะแนนเฉลี่ยของนักเรียนทั้ง 3 กลุ่มแตกต่างกันอย่างมีนัยสำคัญทางสถิติ

ข้อควรระวัง

1. ANOVA บอกได้แค่ ว่า มีความแตกต่างระหว่างอย่างน้อย 1 กลุ่ม แต่ไม่บอก ว่ากลุ่มใดต่าง ต้องใช้ Post-hoc test ต่อ (เช่น Tukey, Bonferroni)

2. ต้องตรวจสอบสมมติฐาน ได้แก่ ความปกติ (Normality) และความแปรปรวนเท่ากัน (Homogeneity of variances)

3. ข้อมูล Outliers มีผลกระทบมาก อาจทำให้ค่า F บิดเบือน

สรุปได้ว่า F -test เป็นเครื่องมือสำคัญใน ANOVA สำหรับเปรียบเทียบค่าเฉลี่ยมากกว่า 2 กลุ่มพร้อมกัน โดยพิจารณาจากสัดส่วนของความแปรปรวนระหว่างกลุ่มและภายในกลุ่ม การตีความต้องดูทั้งค่า F , p -value และทำ Post-hoc test ต่อเพื่อระบุว่าคู่ใดแตกต่างกัน

2) การวิเคราะห์ค่า p (p-value Analysis)

ค่า p -value คือ ความน่าจะเป็น (probability) ที่จะได้ผลการทดสอบทางสถิติ (เช่น ค่า F , t หรือ Z) ที่มีความสุดโต่ง (extreme) เท่ากับหรือมากกว่าค่าที่สังเกตได้จริง ภายใต้สมมติฐานศูนย์ (H_0) ว่ายังคงเป็นจริง (Howell, 2007; Field, 2013) กล่าวให้ง่ายขึ้น: ค่า p บอกว่า “ถ้า H_0 เป็นจริง โอกาสที่จะได้ผลการทดลองแบบนี้ หรือมากกว่านี้มีเท่าไร”

- ค่า p ต่ำมาก (เช่น ≤ 0.05) $\rightarrow$ ผลลัพธ์ไม่น่าจะเกิดขึ้นหาก H_0 จริง $\rightarrow$ มีเหตุผลเพียงพอในการปฏิเสธ H_0

- ค่า p สูง (เช่น > 0.05) $\rightarrow$ ผลลัพธ์ที่ได้ยังสอดคล้องกับ H_0 $\rightarrow$ ไม่มีหลักฐานเพียงพอที่จะปฏิเสธ H_0

ค่า p ในการวิเคราะห์ความแปรปรวน (p-value in ANOVA)

ในการทดสอบ ANOVA นักวิจัยต้องการทราบว่า ค่าเฉลี่ยของกลุ่มต่าง ๆ มีความแตกต่างกันหรือไม่ โดยใช้ค่า F -test ซึ่งเปรียบเทียบ

- ความแปรปรวนระหว่างกลุ่ม ($MS_{between}$) กับ

- ความแปรปรวนภายในกลุ่ม (MS_{within})

ถ้า F ที่ได้มีค่า สูงมาก โอกาสที่จะเกิดขึ้นโดยบังเอิญเมื่อ H_0 เป็นจริงจะน้อย $\rightarrow$ ทำให้ค่า p เล็กตามไปด้วย

สมมติฐานที่ใช้ทดสอบ

- H_0 : ค่าเฉลี่ยทุกกลุ่มเท่ากัน ($\mu_1 = \mu_2 = \mu_3 \dots \mu_k$)
- H_1 : อย่างน้อยหนึ่งกลุ่มมีค่าเฉลี่ยแตกต่าง

หลักเกณฑ์การตัดสินใจ

- ถ้า $p \leq .05 \rightarrow$ ปฏิเสธ $H_0 \rightarrow$ สรุปว่ามีความแตกต่างระหว่างกลุ่มอย่างน้อยหนึ่งคู่
- ถ้า $p > .05 \rightarrow$ ไม่ปฏิเสธ $H_0 \rightarrow$ สรุปว่ายังไม่มีหลักฐานเพียงพอว่ากลุ่มต่างกัน

ตัวอย่างการคำนวณและการตีความค่า p

นักวิจัยต้องการเปรียบเทียบคะแนนสอบวิชาสถิติของนักเรียน 3 กลุ่ม (A, B, C) ผล ANOVA ได้ว่า

- ค่า $F = 43.6$
- $df_{between} = 2, df_{within} = 6$
- $p < 0.001$

การตีความ

เนื่องจาก $p < .05$ จึงปฏิเสธ $H_0 \rightarrow$ ค่าเฉลี่ยคะแนนสอบของนักเรียนทั้งสามกลุ่มแตกต่างกันอย่างมีนัยสำคัญทางสถิติ อย่างไรก็ตาม ยังไม่รู้ว่ากลุ่มใดแตกต่างกัน ต้องใช้การวิเคราะห์แบบ Post-hoc เช่น Tukey HSD ต่อไป

ข้อควรระวังในการใช้ค่า p

1. p ไม่ใช่ความน่าจะเป็นที่ H_0 จะเป็นจริงหรือเท็จ

หลายคนเข้าใจผิดว่าค่า $p = 0.03$ หมายความว่า H_0 มีโอกาส 3% ที่จะเป็นจริง $\rightarrow$ ไม่ถูกต้อง ความหมายที่ถูกต้องคือ “ถ้า H_0 เป็นจริง จะมีโอกาส 3% ที่ได้ผลลัพธ์สุดโต่งเช่นนี้” (Cohen, 1988)

2. ค่า p ไม่ได้บอกขนาดของความแตกต่าง (effect size)

ค่า p เพียงบอกว่ามี ความแตกต่างที่น่าจะ “จริง” ตามเกณฑ์สถิติ แต่ไม่ได้บอกว่าแตกต่างมากหรือน้อย $\rightarrow$ นักวิจัยควรรายงาน effect size เช่น η^2, ω^2 หรือ Cohen's d ควบคู่ไปด้วย

3. ค่า p ขึ้นกับขนาดตัวอย่าง (sample size)

- ตัวอย่างใหญ่มาก $\rightarrow$ แม้ความแตกต่างเล็กน้อยก็อาจได้ p ที่มีนัยสำคัญ
- ตัวอย่างเล็ก $\rightarrow$ แม้ความแตกต่างมากก็อาจไม่พบความนัยสำคัญ

4. การตีความต้องพิจารณาบริบท (contextual interpretation)

ควรอ้างอิงทฤษฎีและงานวิจัยที่เกี่ยวข้องร่วมด้วย ไม่ควรสรุปผลเพียงเพราะ $p < .05$

สรุปได้ว่า การวิเคราะห์ค่า p ใน ANOVA เป็นขั้นตอนสำคัญในการพิจารณาว่าควรปฏิเสธสมมติฐานศูนย์หรือไม่ โดย $p \leq .05$ บ่งชี้ว่าค่าเฉลี่ยของกลุ่มต่าง ๆ มีความแตกต่างกันอย่างมีนัยสำคัญ แต่การตีความที่ถูกต้องควรพิจารณาร่วมกับ ค่า F , effect size และการทดสอบ Post-hoc เพื่อให้ได้ข้อสรุปที่สมบูรณ์ทั้งเชิงสถิติและเชิงปฏิบัติ

3) การวิเคราะห์แบบ Post-hoc (Tukey HSD Test)

เมื่อทำการวิเคราะห์ความแปรปรวน (ANOVA) แล้วพบว่า มีความแตกต่างระหว่างกลุ่มอย่างมีนัยสำคัญ ($p \leq 0.05$) จะยังไม่ทราบว่าความแตกต่างนั้นเกิดขึ้นที่ คู่ใดของกลุ่ม การทดสอบแบบ Post-hoc (Multiple Comparison Tests) จึงถูกนำมาใช้เพื่อตรวจสอบว่า กลุ่มใดต่างจากกลุ่มใด โดยควบคุมความผิดพลาดแบบที่ 1 (Type I Error) จากการเปรียบเทียบหลายครั้ง (Howell, 2007; Field, 2013)

Tukey HSD เป็นวิธีที่นิยมมากที่สุดในการทดสอบความแตกต่างของค่าเฉลี่ยคู่ต่อคู่หลัง ANOVA เพราะ

- ควบคุม Type I Error ได้ดี
- เหมาะกับกลุ่มที่มีขนาดเท่ากัน
- ตีความได้ง่ายตรงไปตรงมา

สูตรการคำนวณ

$$HSD = q_{\alpha, k, df_{within}} \times \sqrt{\frac{MS_{within}}{n}}$$

โดยที่

- $q_{\alpha, k, df_{within}}$ = ค่า critical จากตาราง Studentized Range Distribution
- MS_{within} = Mean Square ภายในกลุ่ม (จากตาราง ANOVA)
- n = จำนวนตัวอย่างต่อกลุ่ม (สมมติเท่ากันทุกกลุ่ม)
- k = จำนวนกลุ่ม

หลักการตัดสินใจ

- ถ้า $|\bar{x}_i - \bar{x}_j| > HSD \rightarrow$ กลุ่ม i และ j แตกต่างกันอย่างมีนัยสำคัญ
- ถ้า $|\bar{x}_i - \bar{x}_j| \leq HSD \rightarrow$ กลุ่ม i และ j ไม่แตกต่างกัน

ตัวอย่างการคำนวณ Tukey HSD

ใช้ข้อมูลจากตัวอย่าง ANOVA (นักเรียน 3 กลุ่ม)

ตารางที่ 7.4 ค่าเฉลี่ยผลสัมฤทธิ์ทางการเรียนของนักเรียน

กลุ่ม	คะแนน	ค่าเฉลี่ย
A	70, 75, 80	75.0
B	65, 68, 72	68.3
C	85, 90, 95	90.0

ผลจาก ANOVA

- $MS_{within} = 7.78$

- $n = 3$ ต่อกลุ่ม

- $k = 3$ กลุ่ม

- $df_{within} = 6$

หาค่า q จากตาราง Studentized Range ($\alpha = 0.05$, $k = 3$, $df = 6$) ≈ 4.34

$$HSD = 4.34 \times \sqrt{\frac{7.78}{3}} = 4.34 \times 1.61 \approx 7.0$$

เปรียบเทียบค่าเฉลี่ยคู่ต่อคู่

- A vs B: $|75 - 68.3| = 6.7 < 7.0 \rightarrow$ ไม่ต่าง

- A vs C: $|75 - 90| = 15 > 7.0 \rightarrow$ ต่าง

- B vs C: $|68.3 - 90| = 21.7 > 7.0 \rightarrow$ ต่าง

สรุป กลุ่ม C (สอนด้วยวิธีใหม่) มีคะแนนสูงกว่ากลุ่ม A และ B อย่างมีนัยสำคัญ แต่ A และ B ไม่แตกต่างกัน

การตีความผล

Tukey HSD ไม่เพียงระบุว่ามีความแตกต่าง แต่ยังบอกได้ชัดเจนว่าคู่ใดแตกต่างกัน การตีความควรเชื่อมโยงกับทฤษฎี เช่น ในการศึกษาด้านการศึกษา ผลที่ได้อาจบ่งชี้ว่าวิธีการสอน C มีประสิทธิภาพสูงกว่าวิธี A และ B

ข้อควรระวัง

1. เหมาะกับกลุ่มที่มีจำนวนสมาชิกเท่ากัน หากไม่เท่าควรใช้วิธีอื่น เช่น Games-Howell Test

2. ผลลัพธ์ควรพิจารณาพร้อมกับ ค่าเฉลี่ยและขนาดอิทธิพล (Effect Size, η^2 หรือ Cohen's d) เพื่อประเมินความสำคัญเชิงปฏิบัติ

3. Tukey ค่อนข้าง อนุรักษ์นิยม (conservative) $\rightarrow$ โอกาสพบความแตกต่างน้อยกว่าวิธี Bonferroni แต่มีความน่าเชื่อถือสูง

สรุปได้ว่า การวิเคราะห์แบบ Post-hoc โดยใช้ Tukey HSD Test เป็นวิธีมาตรฐานที่ช่วยตรวจสอบว่า คู่ของกลุ่มใดมีค่าเฉลี่ยแตกต่างกัน หลังจากที่ได้ผลการวิเคราะห์ ANOVA แสดงว่ามีความแตกต่างโดยรวมอย่างมีนัยสำคัญทางสถิติ ข้อดีคือสามารถควบคุม Type I Error ได้ดี เหมาะกับกรณีที่แต่ละกลุ่มมีจำนวนสมาชิกเท่ากัน และผลที่ได้ตีความได้ง่ายตรงไปตรงมา อย่างไรก็ตาม ควรใช้ร่วมกับตัวชี้วัดอื่น เช่น Effect Size เพื่อประเมินความสำคัญเชิงปฏิบัติ และในกรณีที่กลุ่มมีขนาดไม่เท่ากันหรือความแปรปรวนไม่เท่ากัน ควรเลือกวิธีที่เหมาะสม เช่น Games–Howell Test

5. ข้อกำหนดเบื้องต้นและการตรวจสอบสมมติฐาน

การใช้สถิติเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม เช่น t-test และ ANOVA ไม่สามารถนำไปใช้กับข้อมูลใด ๆ ได้โดยตรง แต่ต้องมีการตรวจสอบ ข้อกำหนดเบื้องต้น (statistical assumptions) เสียก่อน เพื่อให้การวิเคราะห์มีความถูกต้อง (validity) และผลลัพธ์ที่ได้เชื่อถือได้ หากไม่ตรวจสอบสมมติฐาน ผลการทดสอบอาจคลาดเคลื่อน ทำให้สรุปผิดพลาด

1) ตัวแปรตามต้องเป็นเชิงปริมาณ (Interval/Ratio)

- ตัวแปรตาม (Dependent Variable) ต้องวัดในระดับอันตรภาค (Interval) หรืออัตราส่วน (Ratio)

- Interval: มีระยะห่างระหว่างค่าที่เท่ากัน เช่น คะแนนสอบ อุณหภูมิ (°C)

- Ratio: มีศูนย์แท้จริง เช่น รายได้ น้ำหนัก ส่วนสูง เวลา

- **เหตุผล** t-test และ ANOVA ใช้ค่าเฉลี่ยและความแปรปรวน ซึ่งต้องอ้างอิงจากตัวเลขที่มีความหมายในเชิงปริมาณ

- **วิธีตรวจสอบ** ตรวจสอบมาตราของข้อมูลจากแบบสอบถามหรือเครื่องมือวัด ถ้าเป็น Nominal หรือ Ordinal ต้องเลือกสถิติแบบไม่อิงพารามิเตอร์ เช่น Mann–Whitney U Test หรือ Kruskal–Wallis Test

2) กลุ่มต้องเป็นอิสระจากกัน (ในกรณี Independent t-test/ANOVA)

- สมาชิกในแต่ละกลุ่มต้องไม่ซ้ำกัน และการสุ่มเข้ากลุ่มใดกลุ่มหนึ่งไม่ควรส่งผลต่อการสุ่มอีกกลุ่ม

- **เหตุผล** การละเมิดสมมติฐานนี้ทำให้ค่าความแปรปรวนต่ำกว่าความจริง และค่า t หรือ F ที่ได้สูงเกินจริง → เสี่ยงต่อการเกิด Type I Error

- ตัวอย่าง เช่น

- Independent t-test → เปรียบเทียบนักเรียนชาย vs หญิง ต้องไม่มีใครซ้ำกัน

- ANOVA → เปรียบเทียบกลุ่มนักเรียน 3 ห้องเรียน แต่ละคนอยู่ได้เพียงห้องเดียว

- **วิธีตรวจสอบ** ตรวจสอบการออกแบบการวิจัย → ถ้าเป็นข้อมูลที่ซ้ำวัด (เช่น pre-test/post-test) ต้องใช้ Paired t-test หรือ Repeated Measures ANOVA แทน

3) การแจกแจงต้องใกล้เคียงปกติ (Normality)

- ข้อมูลในแต่ละกลุ่มควรมีการแจกแจงใกล้เคียงโค้งปกติ (Normal Distribution) โดยเฉพาะเมื่อขนาดตัวอย่างน้อยกว่า 30

- **เหตุผล** การใช้ t-test และ ANOVA อาศัยทฤษฎีการแจกแจงปกติในการอ้างอิงค่าสถิติ หากข้อมูลเบ้มาก ผลลัพธ์จะไม่น่าเชื่อถือ

- วิธีตรวจสอบ

1) Shapiro–Wilk Test และ Kolmogorov–Smirnov Test → $p > .05$ แสดงว่าไม่ต่างจากปกติ

2) Histogram และ Q–Q Plot → ใช้ดูรูปร่างข้อมูลด้วยสายตา

3) Skewness และ Kurtosis → ค่าอยู่ในช่วง -2 ถึง +2 ถือว่ายอมรับได้ (Gravetter & Wallnau, 2009)

- **ถ้าละเมิดสมมติฐาน** ใช้การแปลงข้อมูล (log, square root) หรือใช้สถิติแบบไม่อิงพารามิเตอร์ เช่น Wilcoxon Test, Kruskal–Wallis Test

4) ความแปรปรวนของแต่ละกลุ่มต้องเท่ากัน (Homoscedasticity)

- ความแปรปรวน (variance) ของแต่ละกลุ่มควรใกล้เคียงกัน

- **เหตุผล** หากความแปรปรวนไม่เท่ากัน การประมาณค่าเฉลี่ยของกลุ่มอาจคลาดเคลื่อน ทำให้ค่า t-test หรือ F-test ไม่ถูกต้อง

- วิธีตรวจสอบ

- Levene's Test (นิยมที่สุด) → ถ้า $p > .05$ ถือว่าความแปรปรวนเท่ากัน

- Hartley's F-max Test และ Bartlett's Test → ใช้ในบางกรณี

- ถ้าละเมิดสมมติฐาน

- ใช้ Welch's t-test หรือ Welch's ANOVA ที่ไม่สมมติความแปรปรวนเท่ากัน

- ใช้สถิติไม่อิงพารามิเตอร์

5) ข้อควรระวังและแนวทางแก้ไข

(1) การออกแบบการวิจัย → สมมติฐานหลายข้อสามารถป้องกันได้ตั้งแต่การออกแบบ เช่น การสุ่มตัวอย่างอย่างแท้จริง

(2) การใช้สถิติไม่อิงพารามิเตอร์ → เป็นทางเลือกเมื่อข้อมูลไม่ผ่านสมมติฐาน เช่น Mann–Whitney U Test, Kruskal–Wallis Test

(3) การรายงานผล → ควรรายงานทั้งผลการตรวจสอบสมมติฐาน (เช่น ผล Shapiro–Wilk, Levene’s Test) และผลการวิเคราะห์หลัก เพื่อให้ผู้อ่านมั่นใจว่าข้อมูลมีความถูกต้อง

(4) การใช้ขนาดตัวอย่างใหญ่ ($n \geq 30$ ต่อกลุ่ม) → ตามทฤษฎีบทขีดจำกัดกลาง (Central Limit Theorem) ทำให้สมมติฐาน Normality มีผลน้อยลง แต่ Homogeneity ยังคงสำคัญ

สรุปได้ว่า การใช้สถิติเปรียบเทียบค่าเฉลี่ย (t-test และ ANOVA) จำเป็นต้องตรวจสอบสมมติฐานเบื้องต้น ได้แก่ (1) ตัวแปรตามเป็นเชิงปริมาณ (2) กลุ่มต้องเป็นอิสระ (3) การแจกแจงใกล้เคียงปกติ และ (4) ความแปรปรวนเท่ากัน หากสมมติฐานเหล่านี้ไม่เป็นจริง ควรใช้วิธีการทางสถิติที่เหมาะสมหรือปรับแก้ข้อมูลก่อน เพื่อหลีกเลี่ยงการสรุปผิดพลาดและทำให้นักวิจัยมีความน่าเชื่อถือ

สรุปท้ายบท

การวิเคราะห์ความแตกต่างระหว่างกลุ่ม เป็นขั้นตอนสำคัญของการวิจัยเชิงปริมาณในสังคมศาสตร์ เพราะช่วยตอบคำถามเชิงเปรียบเทียบ เช่น กลุ่มตัวอย่างสองกลุ่มหรือมากกว่านั้นมีค่าเฉลี่ยแตกต่างกันอย่างมีนัยสำคัญทางสถิติหรือไม่ เครื่องมือหลัก ได้แก่ Independent Samples t-test สำหรับการเปรียบเทียบสองกลุ่มอิสระ, Paired Samples t-test สำหรับกลุ่มที่มีการวัดซ้ำหรือจับคู่ และ ANOVA (Analysis of Variance) สำหรับกรณีมีกลุ่มตั้งแต่สามกลุ่มขึ้นไป โดย ANOVA ใช้การทดสอบ F เพื่อวัดสัดส่วนระหว่างความแปรปรวนระหว่างกลุ่มและภายในกลุ่ม เมื่อค่า $p \leq .05$ จึงสรุปว่ามีความแตกต่างอย่างน้อยหนึ่งคู่ และต้องใช้การวิเคราะห์ Post-hoc (เช่น Tukey HSD Test) เพื่อตรวจสอบว่าความแตกต่างเกิดขึ้นที่คู่ใด

อย่างไรก็ตาม การตีความผลต้องอาศัย ข้อกำหนดเบื้องต้น เช่น ตัวแปรตามต้องเป็นเชิงปริมาณ กลุ่มต้องอิสระต่อกัน ข้อมูลควรมีการแจกแจงใกล้เคียงปกติ และความแปรปรวนควรใกล้เคียงกัน หากไม่เป็นไปตามสมมติฐาน ควรเลือกใช้วิธีการทางสถิติที่เหมาะสม เช่น Welch’s t-test, Welch’s ANOVA หรือสถิติไม่อิงพารามิเตอร์ (non-parametric tests) นอกจากนี้ การสรุปผลไม่ควรอาศัยค่า p-value เพียงอย่างเดียว แต่ควรรายงานร่วมกับ Effect Size เพื่อสะท้อนความสำคัญเชิงปฏิบัติ ทั้งหมดนี้ทำให้การเปรียบเทียบค่าเฉลี่ยระหว่างกลุ่มมีความแม่นยำและเชื่อถือได้ สามารถนำไปใช้ในการอธิบายปรากฏการณ์และสนับสนุนการกำหนดนโยบายทางสังคมได้อย่างมีประสิทธิภาพ

บทที่ 8

การใช้โปรแกรมทางสถิติในการวิเคราะห์ข้อมูล (USING STATISTICAL PROGRAMS TO ANALYZE DATA)

ในการทำวิจัยเชิงปริมาณ การวิเคราะห์ข้อมูลทางสถิติถือเป็นขั้นตอนสำคัญที่ช่วยให้ผู้วิจัยสามารถตอบคำถามวิจัย ทดสอบสมมติฐาน และสร้างข้อสรุปเชิงวิชาการได้อย่างเป็นระบบ เดิมทีการวิเคราะห์ข้อมูลต้องอาศัยการคำนวณด้วยมือ ซึ่งเสี่ยงต่อความผิดพลาดและใช้เวลานาน แต่ด้วยความก้าวหน้าของเทคโนโลยีสารสนเทศ ได้มีการพัฒนาโปรแกรมคอมพิวเตอร์เพื่อช่วยให้นักวิจัยสามารถวิเคราะห์ข้อมูลได้อย่างรวดเร็ว ถูกต้อง และมีความน่าเชื่อถือ

โปรแกรมทางสถิติ เช่น SPSS, JASP, R และ PSPP ได้รับความนิยมอย่างแพร่หลาย โดยแต่ละโปรแกรมมีจุดเด่นและข้อจำกัดที่ต่างกัน SPSS เน้นความสะดวกและครอบคลุม, JASP ฟรีและรองรับ Bayesian Analysis, R มีความยืดหยุ่นและทรงพลังในงานขั้นสูง ส่วน PSPP เป็นทางเลือกฟรีที่ใกล้เคียง SPSS เหมาะกับนักศึกษาและผู้วิจัยทั่วไป

ดังนั้น การทำความเข้าใจคุณสมบัติ วิธีใช้งาน และข้อควรระวังในการใช้โปรแกรมเหล่านี้ ไม่เพียงช่วยให้ผู้วิจัยเลือกเครื่องมือที่เหมาะสม แต่ยังทำให้การวิเคราะห์ข้อมูลมีความถูกต้อง มีมาตรฐาน และสามารถนำเสนอผลลัพธ์ในเชิงวิชาการได้อย่างน่าเชื่อถือ

1. แนะนำโปรแกรม SPSS, JASP, R และ PSPP

การวิเคราะห์ข้อมูลทางสถิติในปัจจุบันมีโปรแกรมคอมพิวเตอร์หลากหลายชนิดที่ช่วยให้การคำนวณสะดวก รวดเร็ว และลดข้อผิดพลาดที่อาจเกิดจากการคำนวณด้วยมือ โปรแกรมที่นิยมใช้ในงานวิจัยด้านสังคมศาสตร์ การศึกษา และรัฐประศาสนศาสตร์ ได้แก่ SPSS, JASP, R และ PSPP แต่ละโปรแกรมมีจุดเด่นและข้อจำกัดที่แตกต่างกัน ดังนี้

1) โปรแกรม SPSS (Statistical Package for the Social Sciences)

(1) ความเป็นมาและความสำคัญ

โปรแกรม SPSS ได้รับการพัฒนาขึ้นครั้งแรกในปี ค.ศ. 1968 โดย Norman H. Nie, C. Hadlai Hull และ Dale H. Bent เพื่อใช้ในการวิเคราะห์ข้อมูลด้านสังคมศาสตร์ ก่อนที่จะถูก IBM เข้าซื้อกิจการในปี ค.ศ. 2009 และเปลี่ยนชื่อเป็น IBM SPSS Statistics

โปรแกรมนี้จึงกลายเป็นซอฟต์แวร์มาตรฐานในการทำงานวิจัยเชิงปริมาณ ทั้งในแวดวงสังคมศาสตร์ จิตวิทยา การศึกษา บริหารธุรกิจ และรัฐประศาสนศาสตร์ (Field, 2018)

(2) คุณสมบัติและจุดเด่น

- อินเทอร์เฟซใช้งานง่าย (User-friendly Interface) ผู้ใช้สามารถป้อนข้อมูล วิเคราะห์ และอ่านผลได้โดยไม่ต้องเขียนโค้ด
- เครื่องมือวิเคราะห์ครบถ้วน ตั้งแต่สถิติพื้นฐาน เช่น ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน ไปจนถึงสถิติขั้นสูง เช่น Multiple Regression, Logistic Regression, Factor Analysis, Cluster Analysis
- รองรับข้อมูลหลากหลาย สามารถนำเข้าข้อมูลจาก Excel, CSV, หรือฐานข้อมูล SQL
- ระบบ Output ที่ชัดเจน SPSS แยกหน้าต่างผลลัพธ์ออกจากหน้าต่างข้อมูล ทำให้อ่านง่าย สามารถส่งออกเป็น Word, PDF, Excel ได้

(3) ข้อจำกัด

- ค่าใช้จ่ายสูง ต้องซื้อไลเซนส์ จึงไม่เหมาะกับผู้ใช้ทั่วไปที่มีงบจำกัด
- ความยืดหยุ่นต่ำกว่า R ไม่สามารถปรับแต่งการวิเคราะห์ได้ซับซ้อนเท่ากับภาษา R
- ใช้ทรัพยากรเครื่องสูง ไฟล์ข้อมูลขนาดใหญ่ทำให้โปรแกรมทำงานช้า

(4) การใช้งานในเชิงวิจัย

SPSS เหมาะสำหรับงานวิจัยที่มีข้อมูลจากแบบสอบถาม (Survey Research) และต้องการการวิเคราะห์เชิงสถิติ เช่น

- สถิติเชิงพรรณนา (Descriptive Statistics) ใช้หาค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และตารางความถี่
- การทดสอบสมมติฐาน (Hypothesis Testing) เช่น Independent t-test เพื่อเปรียบเทียบค่าเฉลี่ยของสองกลุ่ม
- การวิเคราะห์ความสัมพันธ์ (Correlation) เช่น Pearson's r เพื่อวัดความสัมพันธ์ระหว่างตัวแปร
- การวิเคราะห์การถดถอย (Regression Analysis): เพื่อทำนายค่าตัวแปรตามจากตัวแปรอิสระ
- การวิเคราะห์ความแปรปรวน (ANOVA): เพื่อเปรียบเทียบค่าเฉลี่ยของหลายกลุ่ม

(5) ตัวอย่างการใช้งานจริง

หากนักวิจัยต้องการตรวจสอบว่า นักเรียนชายและหญิงมีผลสัมฤทธิ์ทางการเรียนแตกต่างกันหรือไม่ สามารถใช้ SPSS ทำ Independent Samples t-test ได้ทันที โดยป้อนข้อมูลตัวแปรเพศและคะแนนสอบ จากนั้นเลือกเมนู *Analyze* → *Compare Means* → *Independent-Samples T Test* ผลลัพธ์ที่ได้จะแสดงค่า t , df และ p -value ซึ่งช่วยให้ผู้วิจัยตัดสินใจได้ว่าความแตกต่างนั้นมีนัยสำคัญทางสถิติหรือไม่

(6) ภาพตัวอย่างโปรแกรม SPSS

	VAR00001	VAR00002	VAR0003	VAR0004	VAR0005
1	1	9	2	6	
2	4	5	1	6	
3	8	4	4	3	
4	7	2	6	5	
5	6	2	3	6	
6	5	3	5	6	
7	4	2	6	6	
8	10	2	6	6	
10					
11					
12					
13					
14					
15					
16					
17					

แผนภาพที่ 8.1 โปรแกรม SPSS

สรุปได้ว่า SPSS เป็นเครื่องมือสำคัญที่ช่วยให้นักวิจัยสามารถทำงานกับข้อมูลได้อย่างสะดวกและเป็นระบบ แม้จะมีค่าใช้จ่ายสูง แต่ด้วยอินเทอร์เฟซที่ใช้งานง่ายและฟังก์ชันการวิเคราะห์ที่ครอบคลุม ทำให้ SPSS ยังคงเป็นตัวเลือกหลักของงานวิจัยเชิงปริมาณในหลายสาขา โดยเฉพาะสังคมศาสตร์และรัฐประศาสนศาสตร์

2) โปรแกรม JASP (Jeffreys's Amazing Statistics Program)

(1) ความเป็นมาและความสำคัญ

JASP เป็นซอฟต์แวร์โอเพนซอร์สที่พัฒนาขึ้นโดยทีมงานจากมหาวิทยาลัยอัมสเตอร์ดัม (University of Amsterdam) นำโดย Eric-Jan Wagenmakers เพื่อให้เป็น “ทางเลือกใหม่” สำหรับนักวิจัยที่ต้องการโปรแกรมวิเคราะห์สถิติที่ใช้งานง่าย ฟรี และโปร่งใส (Love et al., 2019) จุดเด่นที่ทำให้ JASP แตกต่างคือ การสนับสนุนทั้ง Frequentist Statistics (สถิติแบบดั้งเดิม) และ Bayesian Statistics (สถิติแบบเบย์เซียน) ภายในโปรแกรมเดียว

(2) คุณสมบัติและจุดเด่น

- ใช้งานฟรี (Freeware) เปิดให้ดาวน์โหลดและใช้งานได้โดยไม่เสียค่าใช้จ่าย
- อินเทอร์เฟซที่ใช้งานง่าย หน้าตาคล้าย SPSS แต่เรียบง่ายกว่า เหมาะกับผู้เริ่มต้น
- การแสดงผลทันที (Instant Output) ผลลัพธ์ที่ได้เชื่อมโยงกับตารางและกราฟแบบ Real-time หากผู้ใช้แก้ไขค่าตัวแปร ผลลัพธ์จะอัปเดตอัตโนมัติ
- รองรับ Bayesian Analysis ซึ่งโปรแกรม SPSS ไม่ได้รองรับโดยตรง
- การสร้างกราฟสวยงาม เช่น Boxplots, Histograms, Scatterplots พร้อมจัดรูปแบบอัตโนมัติ

(3) ข้อจำกัด

- ฟังก์ชันยังไม่ครอบคลุม การวิเคราะห์เชิงซับซ้อน เช่น SEM (Structural Equation Modeling) ยังไม่รองรับเต็มรูปแบบ
- ต้องใช้การนำเข้าข้อมูล JASP ไม่เน้นการป้อนข้อมูลโดยตรง แต่เหมาะกับการนำเข้าข้อมูลจาก Excel, CSV
- ฟังก์ชันการอัปเดต แม้จะมีการพัฒนาอย่างต่อเนื่อง แต่บางฟังก์ชันยังอยู่ในช่วงทดลอง

(4) การใช้งานในเชิงวิจัย

JASP เหมาะกับงานวิจัยทางการศึกษาและสังคมศาสตร์ที่ต้องการการวิเคราะห์เชิงพรรณนาและการทดสอบสมมติฐาน ตัวอย่างการใช้งานที่พบบ่อย ได้แก่

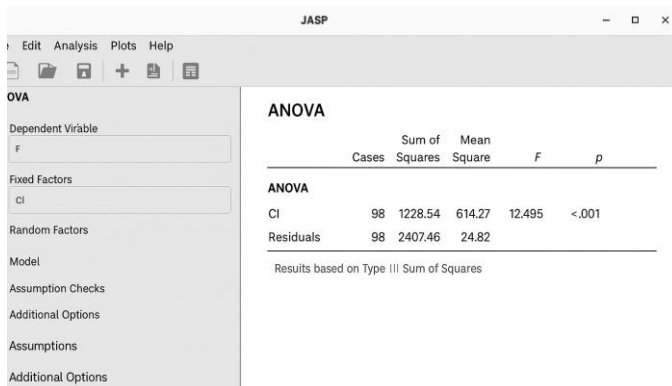
- การวิเคราะห์เชิงพรรณนา (Descriptive Statistics) ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน กราฟความถี่
- การทดสอบ t-test และ ANOVA เปรียบเทียบค่าเฉลี่ยระหว่างกลุ่ม
- Correlation Analysis วิเคราะห์ความสัมพันธ์ระหว่างตัวแปร
- Regression Analysis ทั้งแบบเชิงเส้น (Linear) และเชิงพหุคูณ (Multiple)

- Bayesian Analysis: เช่น Bayesian t-test, Bayesian ANOVA

(5) ตัวอย่างการใช้งานจริง

นักวิจัยที่ต้องการเปรียบเทียบคะแนนสอบระหว่างนักเรียน 3 โรงเรียน สามารถนำข้อมูลเข้า JASP และใช้ One-way ANOVA เพื่อตรวจสอบความแตกต่างค่าเฉลี่ย โปรแกรมจะสร้างตารางสรุปผล F-test และ p-value รวมทั้งกราฟ Boxplot ที่สวยงามโดยอัตโนมัติ

(6) ภาพตัวอย่างโปรแกรม JASP



แผนภาพที่ 8.2 โปรแกรม JASP

สรุปได้ว่า JASP เป็นโปรแกรมทางเลือกที่เหมาะสมสำหรับผู้ที่ต้องการวิเคราะห์ข้อมูลทางสถิติอย่างง่ายและรวดเร็ว โดยไม่ต้องเสียค่าใช้จ่าย จุดเด่นคือการรองรับการวิเคราะห์แบบ Bayesian ซึ่งช่วยเพิ่มมิติใหม่ให้กับงานวิจัย แม้ยังมีข้อจำกัดในบางฟังก์ชันขั้นสูง แต่ก็ถือว่าเป็นโปรแกรมที่น่าจับตามองและได้รับความนิยมเพิ่มขึ้นเรื่อย ๆ

3) โปรแกรม R (และ RStudio)

(1) ความเป็นมาและความสำคัญ

R เป็นภาษาโปรแกรมสำหรับการคำนวณทางสถิติและการสร้างกราฟเชิงวิชาการ พัฒนาขึ้นครั้งแรกในปี ค.ศ. 1993 โดย Ross Ihaka และ Robert Gentleman จากมหาวิทยาลัยโอ๊คแลนด์ (University of Auckland) ประเทศนิวซีแลนด์ ต่อมา มีการพัฒนาต่อเนื่องโดย *R Core Team* ทำให้ R กลายเป็นซอฟต์แวร์โอเพนซอร์สที่ทรงพลังที่สุดด้านสถิติและ Data Science ในปัจจุบัน (R Core Team, 2023)

เพื่อให้ใช้งานได้ง่ายขึ้น มีการพัฒนา RStudio ซึ่งเป็น Integrated Development Environment (IDE) ที่รวมฟังก์ชันต่าง ๆ ไว้ในหน้าจอเดียว ทั้ง Script Editor, Console, Environment/History และ Plot/Packages/Help

(2) คุณสมบัติและจุดเด่น

- ฟรีและโอเพนซอร์ส ดาวน์โหลดและใช้งานได้โดยไม่เสียค่าใช้จ่าย
- รองรับแพ็คเกจจำนวนมาก มีผู้พัฒนาแพ็คเกจ (Packages) มากกว่า 18,000 ตัว เช่น

- ggplot2 สำหรับสร้างกราฟ
- dplyr สำหรับจัดการข้อมูล
- caret สำหรับ Machine Learning

ธุรกิจ

- รองรับงาน Big Data และ AI ใช้ได้ทั้งงานวิจัยเชิงวิชาการและงานเชิง

ละเอียดมาก

- ความยืดหยุ่นสูง สามารถปรับแต่งการวิเคราะห์และการสร้างกราฟได้

- ชุมชนผู้ใช้ใหญ่ มีเอกสาร ตัวอย่าง และฟอรัมสนับสนุนจำนวนมาก

(3) ข้อจำกัด

- ต้องใช้การเขียนโค้ด ผู้ใช้ใหม่อาจรู้สึกว่ายากในการเริ่มต้น

- เรียนรู้โครงสร้างข้อมูลของ R เช่น Vectors, Data Frames, Lists ซึ่งต้องเข้าใจเพื่อใช้งานได้เต็มที่

- อินเทอร์เฟซดั้งเดิมไม่เป็นมิตร หากไม่ใช่ RStudio อาจใช้งานลำบาก

4) การใช้งานในเชิงวิจัย

R และ RStudio เหมาะกับการวิเคราะห์ข้อมูลเชิงซับซ้อนและข้อมูลขนาดใหญ่ ตัวอย่างการใช้งาน ได้แก่

- สถิติพื้นฐาน ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน การทดสอบ t-test

- การวิเคราะห์ขั้นสูง Regression, ANOVA, Factor Analysis, SEM

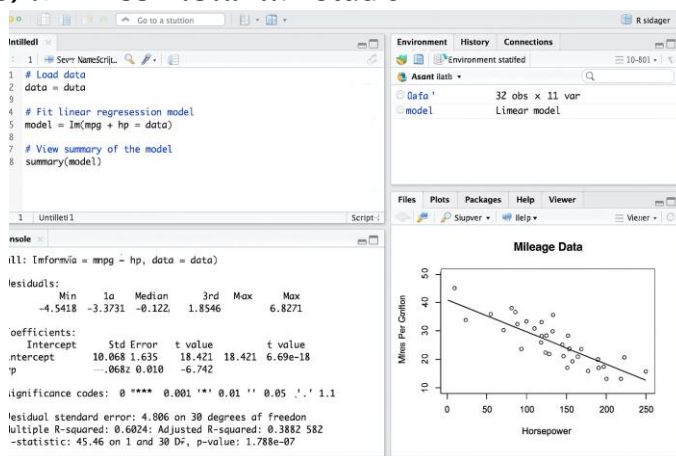
- การทำ Machine Learning Classification, Clustering, Prediction

- การทำ Data Visualization กราฟเชิงวิชาการที่ปรับแต่งได้ละเอียด เช่น

Heatmaps, Interactive Charts

- การเขียนรายงานอัตโนมัติด้วย R Markdown ที่สามารถสร้างไฟล์ Word, PDF, HTML ได้

(5) ภาพตัวอย่างโปรแกรม RStudio



แผนภาพที่ 8.3 โปรแกรม RStudio

สรุปได้ว่า R และ RStudio เป็นเครื่องมือที่ทรงพลังและยืดหยุ่นที่สุดในการวิเคราะห์ข้อมูลทางสถิติและวิทยาศาสตร์ข้อมูล แม้จะมีความยากในช่วงเริ่มต้นเพราะต้องเขียนโค้ด แต่ด้วยคุณสมบัติฟรี ครอบคลุมทุกการวิเคราะห์ และชุมชนผู้ใช้ที่กว้างขวาง ทำให้ R กลายเป็นโปรแกรมที่นักวิจัยและนักวิเคราะห์ข้อมูลทั่วโลกนิยมใช้

4) โปรแกรม PSPP

(1) ความเป็นมาและความสำคัญ

PSPP เป็นโปรแกรมโอเพนซอร์สที่พัฒนาโดย *Free Software Foundation (FSF)* ภายใต้โครงการ GNU โดยมีจุดประสงค์เพื่อเป็นทางเลือกแทน SPSS ซึ่งเป็นซอฟต์แวร์เชิงพาณิชย์ (GNU PSPP, 2022) PSPP ถูกออกแบบให้ใช้งานง่าย มีหน้าต่างและฟังก์ชันคล้าย SPSS เพื่อช่วยให้นักวิจัยและนักศึกษา โดยเฉพาะในประเทศกำลังพัฒนา สามารถเข้าถึงเครื่องมือวิเคราะห์สถิติได้โดยไม่ต้องเสียค่าใช้จ่ายสูง

(2) คุณสมบัติและจุดเด่น

- ฟรีและโอเพนซอร์ส ใช้งานได้โดยไม่เสียค่าใช้จ่าย
- อินเทอร์เฟซคล้าย SPSS มีทั้ง Data View และ Variable View ทำให้ผู้ที่เคยใช้ SPSS มาก่อนปรับตัวได้ทันที
- การทำงานที่รวดเร็ว เหมาะกับการวิเคราะห์ข้อมูลปริมาณมากในเวลาอันสั้น
- รองรับหลายแพลตฟอร์ม ติดตั้งได้ทั้ง Windows, macOS และ Linux
- เหมาะกับงานวิจัยเชิงพรรณนาและทดสอบสมมติฐาน เช่น t-test, ANOVA, Regression, Correlation

(3) ข้อจำกัด

- ฟังก์ชันไม่ครอบคลุม ยังไม่รองรับการวิเคราะห์ขั้นสูง เช่น SEM (Structural Equation Modeling) หรือ Multilevel Analysis
- การแสดงผล Output ไม่ยืดหยุ่นเท่า SPSS ตารางและกราฟอาจต้องนำไปปรับแต่งเพิ่มเติม
- การพัฒนาอัปเดตช้า เมื่อเทียบกับซอฟต์แวร์อื่นที่มีการพัฒนาอย่างต่อเนื่อง

(4) การใช้งานในเชิงวิจัย

PSPP เหมาะสำหรับงานวิจัยเชิงปริมาณทั่วไป เช่น

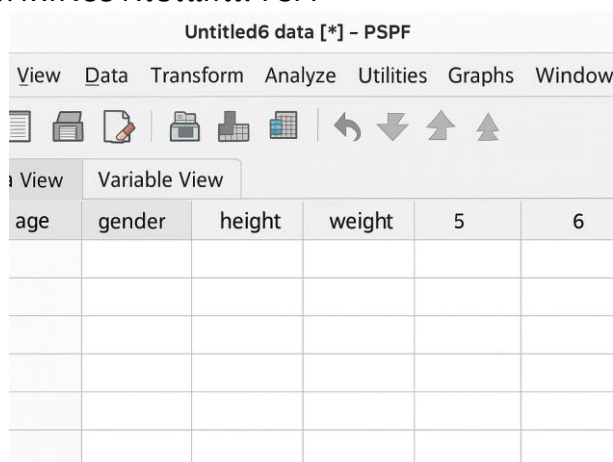
- Descriptive Statistics ค่ากลาง ค่าการกระจาย ตารางความถี่
- Hypothesis Testing Independent t-test, Paired t-test

- Correlation Analysis Pearson และ Spearman
- Regression Analysis Linear Regression
- ANOVA การเปรียบเทียบค่าเฉลี่ยหลายกลุ่ม

(5) ตัวอย่างการใช้งานจริง

สมมติว่านักวิจัยต้องการทดสอบว่านักศึกษาปริญญาตรีและปริญญาโทมีความพึงพอใจในการเรียนการสอนแตกต่างกันหรือไม่ นักวิจัยสามารถป้อนข้อมูลลงใน **Data View** ของ SPSS และเลือกเมนู *Analyze* → *Compare Means* → *Independent-Samples T Test* โปรแกรมจะคำนวณค่า t-test และแสดงค่า p-value ช่วยตัดสินผลการวิจัยได้ทันที

(6) ภาพตัวอย่างโปรแกรม SPSS



แผนภาพที่ 8.4 โปรแกรม SPSS

สรุปได้ว่า SPSS เป็นซอฟต์แวร์ที่ออกแบบมาเพื่อให้ผู้ใช้เข้าถึงการวิเคราะห์สถิติได้โดยไม่ต้องพึ่งพาโปรแกรมเชิงพาณิชย์ที่มีค่าใช้จ่ายสูง แม้ยังมีข้อจำกัดในฟังก์ชันขั้นสูง แต่ก็เหมาะกับงานวิจัยพื้นฐานและการเรียนการสอน โดยเฉพาะในกลุ่มนักศึกษาและนักวิจัยที่ต้องการใช้ซอฟต์แวร์ฟรี

2. การป้อนข้อมูลและการจัดโครงสร้าง Dataset

การป้อนข้อมูล (Data Entry) และการจัดโครงสร้างชุดข้อมูล (Dataset Structure) เป็น “รากฐาน” ของงานวิเคราะห์ทางสถิติ หากข้อมูลมีข้อผิดพลาด เช่น ป้อนตัวเลขผิด โครงสร้างข้อมูลไม่เป็นระบบ หรือไม่จัดการ Missing values อย่างเหมาะสม จะส่งผลให้ผลการวิเคราะห์บิดเบือนทันที (Creswell & Creswell, 2018) ดังนั้น งานวิจัยที่ดีต้องให้ความสำคัญกับขั้นตอนนี้ไม่แพ้การเลือกสถิติหรือการตีความผล

1) การกำหนดตัวแปร (Variables)

- หลักการตั้งชื่อ (Naming conventions)

- ใช้ชื่อสั้น กระชับ และสื่อความหมาย เช่น age, gender, income
- หลีกเลี่ยงการเว้นวรรค ใช้ขีดกลาง (_) แทน เช่น score_pre
- หลีกเลี่ยงอักขรพิเศษ เช่น %, &, @
- แนะนำให้ใช้ตัวอักษรภาษาอังกฤษ แม้ข้อมูลเป็นภาษาไทย (เพื่อหลีกเลี่ยงปัญหาเข้ารหัส)

- ประเภทของตัวแปร (Measurement Scales)

1. Nominal (นามบัญญัติ)

- จัดกลุ่ม ไม่มีลำดับ เช่น เพศ (1 = ชาย, 2 = หญิง)

2. Ordinal (ลำดับ)

- มีลำดับ แต่ช่วงห่างไม่เท่ากัน เช่น ระดับความพึงพอใจ (1 = น้อยมาก ... 5 = มากที่สุด)

3. Interval/Ratio (อันตรภาค/อัตราส่วน)

- มีค่าต่อเนื่องและช่วงห่างเท่ากัน
- Interval → ไม่มีศูนย์แท้ เช่น อุณหภูมิ (°C)
- Ratio → มีศูนย์แท้ เช่น อายุ รายได้ คะแนนสอบ

2) การป้อนข้อมูล (Data Entry)

- SPSS/PSPP

- ใช้ **Variable View** → กำหนดตัวแปร (Name, Label, Type, Measure, Value labels, Missing)
- ใช้ **Data View** → กรอกข้อมูลเป็นแถว (rows = cases, columns = variables)

- ตัวอย่าง Variable View

ตารางที่ 8.1 Variable View

Name	Label	Values	Missing	Measure
gender	เพศ	1=ชาย, 2=หญิง	9	Nominal
age	อายุ (ปี)	-	99	Scale
edu	การศึกษา	1=HS, 2=BA, 3=Grad	-	Ordinal

- Syntax ตัวอย่าง (นำเข้าจาก CSV)

```
GET DATA /TYPE=TEXT
  /FILE="C:\data\dataset.csv"
  /DELCASE=LINE
  /DELIMITERS=","
  /FIRSTCASE=2
  /VARIABLES= id F8.0 gender F1.0 age F3.0 edu F1.0
  score_pre F3.0 score_post F3.0.
EXECUTE.
```

- JASP

- นำเข้าข้อมูลจาก Excel/CSV ได้ทันที
- ระบบจะตรวจจับตัวแปรเชิงตัวเลข/ข้อความอัตโนมัติ
- ผู้ใช้สามารถปรับประเภทเป็น Nominal/ Ordinal/ Scale ได้จากเมนู

- ขั้นตอน

1. File → Open → Computer → Browse
2. เลือกไฟล์ Excel/CSV
3. ตรวจสอบและแก้ไขประเภทตัวแปรในเมนู

- R

- ป้อนข้อมูลด้วย Data frame หรือ import จากไฟล์ CSV/Excel

สร้าง Data frame ด้วยมือ

```
data <- data.frame(
  id = c(1,2,3),
  gender = c(1,2,1),
  age = c(20,22,21),
  score_pre = c(50,55,60),
  score_post = c(65,70,75)
)
```

Import CSV

```
data <- read.csv("data.csv", header=TRUE)
```

กำหนดค่า Factor และ Missing

```
data$gender <- factor(data$gender, levels=c(1,2), labels=c("ชาย", "หญิง"))
```

```
data$age[data$age==99] <- NA
```

3) การตรวจสอบข้อมูล (Data Validation)

- Range check

- ตรวจสอบค่าที่อยู่นอกช่วงสมเหตุสมผล เช่น อายุ 300 ปี

- Logic check

- ตรวจสอบความสัมพันธ์ เช่น อายุ 12 ปี แต่การศึกษา = ปริญญาโท

- Outliers

- ตรวจสอบด้วย Z-score ($> \pm 3$) หรือ Boxplot

```
boxplot(data$score_post, main="Boxplot of Post-test Score")
```

- Missing values

- ตรวจสอบจำนวนค่าที่หายไปและรูปแบบการหาย (MCAR, MAR, MNAR)

4) การทำ Codebook

- Codebook คือ “คู่มือ Dataset” ช่วยกำหนดมาตรฐาน ลดความสับสน และทำให้ข้อมูลใช้ซ้ำได้

ตัวอย่าง Codebook

ตารางที่ 8.2 Codebook

Variable	Label	Type/Measure	Valid range	Missing	Notes
id	รหัสผู้ตอบ	Scale	001-999	-	Key
gender	เพศ	Nominal	1-2	9	1=ชาย, 2=หญิง
age	อายุ (ปี)	Scale	15-80	99	ตรวจสอบ Outlier
score_pre	คะแนนก่อนเรียน	Scale	0-100	-	ใช้ paired t-test
score_post	คะแนนหลังเรียน	Scale	0-100	-	ใช้ Independent t-test

สรุปได้ว่า การป้อนข้อมูลและการจัดโครงสร้าง Dataset เป็นขั้นตอนที่วางรากฐานให้งานวิจัยมีคุณภาพ ตั้งแต่การกำหนดตัวแปร การป้อนข้อมูลในโปรแกรมต่าง ๆ การตรวจสอบความถูกต้อง ไปจนถึงการทำ Codebook หากดำเนินการอย่างเป็นระบบจะช่วยให้การวิเคราะห์ในขั้นต่อไปมีความถูกต้อง น่าเชื่อถือ และสอดคล้องกับมาตรฐานสากล

3. การวิเคราะห์ข้อมูลเชิงพรรณนา (Descriptive Statistics)

การวิเคราะห์ข้อมูลเชิงพรรณนา คือการสรุปคุณลักษณะของข้อมูลในเชิงสถิติพื้นฐาน เพื่อให้เข้าใจภาพรวมของข้อมูลก่อนเข้าสู่การวิเคราะห์เชิงอนุมาน วิธีนี้ช่วยตรวจสอบความสมเหตุสมผลของข้อมูล ค้นหาความผิดปกติ และอธิบายลักษณะสำคัญของกลุ่มตัวอย่าง (Field, 2018)

1) ค่ากลาง (Measures of Central Tendency)

- ค่าเฉลี่ย (Mean) → ใช้เมื่อข้อมูลกระจายปกติ
- มัธยฐาน (Median) → ใช้เมื่อข้อมูลมี outlier
- ฐานนิยม (Mode) → ใช้กับข้อมูลเชิงกลุ่ม เช่น เพศ อาชีพ
- ตัวอย่าง Output SPSS

ตารางที่ 8.3 Output SPSS ค่ากลาง

Gender	N	Mean Score	Std. Deviation
Male	50	42.35	6.12
Female	52	39.88	7.02

2) การวัดการกระจาย (Measures of Dispersion)

- ค่าส่วนเบี่ยงเบนมาตรฐาน (SD) บอกการกระจายรอบค่าเฉลี่ย
- ความแปรปรวน (Variance) ใช้ในสถิติเชิงอนุมาน เช่น ANOVA
- ค่าต่ำสุด-สูงสุด (Range) ตรวจสอบความกว้างของข้อมูล
- ตัวอย่างใน R

```
summary(data$score_post)
```

```
sd(data$score_post, na.rm=TRUE)
```

3) ตารางความถี่และร้อยละ (Frequency Distribution)

- ใช้กับข้อมูลเชิงกลุ่ม เช่น เพศ ระดับการศึกษา
- ตัวอย่าง Output SPSS

ตารางที่ 8.4 Output SPSS ความถี่และร้อยละ

Education Level	Frequency	Percent
High school	40	33.3
Bachelor	60	50.0
Graduate	20	16.7

4) การนำเสนอด้วยกราฟ

- **Histogram** แสดงการแจกแจงของตัวแปรเชิงปริมาณ เช่น คะแนนสอบ

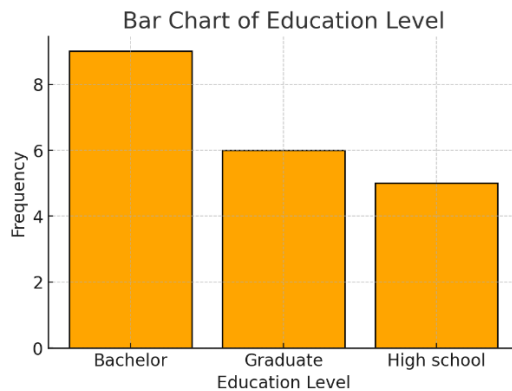
```
hist(data$score_post, main="Histogram of Post-test Scores",  
      xlab="Score", col="lightblue", breaks=10)
```



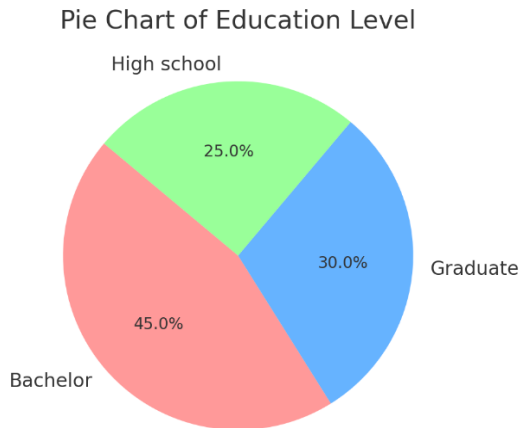
แผนภาพที่ 8.5 Histogram

- **Bar chart, Pie chart** แสดงข้อมูลเชิงกลุ่ม เช่น เพศ ระดับการศึกษา

```
barplot(table(data$edu), col="orange",  
         main="Bar Chart of Education Level")
```



แผนภาพที่ 8.6 Bar chart



แผนภาพที่ 8.7 Pie chart

5) ตัวอย่างตารางสรุปค่าเชิงพรรณนา

ตารางที่ 8.5 Descriptive Statistics

Variable	count	mean	std	min	25%	50%	75%	max	missing
age	120.0	21.86	3.11	18.0	19.0	21.0	24.0	30.0	0
hours	120.0	10.18	3.35	2.0	8.11	9.93	12.25	19.42	0
score_pre	120.0	56.36	10.39	32.1	49.53	56.16	62.27	82.07	0
score_post	119.0	40.26	8.13	25.0	35.02	39.18	45.71	70.64	1

ตารางนี้แสดงค่าเฉลี่ย (mean) ส่วนเบี่ยงเบนมาตรฐาน (std) ค่าต่ำสุด-สูงสุด และจำนวนค่าที่หายไป (missing)

จะเห็นว่า ตัวแปร score_post มีค่า missing = 1 แสดงว่ามีผู้ตอบแบบสอบถามที่ไม่ได้กรอก

สรุปได้ว่า การวิเคราะห์เชิงพรรณนาเป็นขั้นตอนที่ช่วยให้นักวิจัยเข้าใจข้อมูลเบื้องต้น โดยการคำนวณค่ากลาง ค่าการกระจาย การทำตารางความถี่ และการสร้างกราฟ การใช้สถิติพรรณนาไม่เพียงช่วยตรวจสอบคุณภาพข้อมูล แต่ยังเป็นจุดเริ่มต้นที่สำคัญในการตัดสินใจเลือกสถิติอนุมานที่เหมาะสมในขั้นตอนถัดไป

4. การวิเคราะห์ข้อมูลเชิงอนุมาน (Inferential Statistics)

สถิติอนุมาน คือกระบวนการที่ใช้ข้อมูลจากกลุ่มตัวอย่าง (Sample) มาสรุปและตีความเพื่ออ้างอิงไปยังประชากร (Population) ที่ใหญ่กว่า โดยอาศัยทฤษฎีความน่าจะเป็น (Probability Theory) และแบบจำลองทางสถิติ (Statistical Models) (Creswell & Creswell, 2018) การวิเคราะห์ประเภทนี้ช่วยให้ผู้วิจัยสามารถทดสอบสมมติฐาน (Hypothesis Testing) และประมาณค่าพารามิเตอร์ของประชากรได้อย่างมีหลักเกณฑ์

1) ความสำคัญของสถิติอนุมาน

- การทดสอบสมมติฐาน (Hypothesis Testing) ใช้ตรวจสอบว่าข้อสมมติฐานของการวิจัยสอดคล้องกับข้อมูลจริงหรือไม่ เช่น การเปรียบเทียบผลสัมฤทธิ์ของนักเรียนระหว่างกลุ่มทดลองกับกลุ่มควบคุม
- การหาความสัมพันธ์ (Relationship) ศึกษาความสัมพันธ์ของตัวแปร เช่น ความถี่ในการอ่านหนังสือกับผลการสอบ
- การทำนาย (Prediction/Forecasting) ใช้โมเดลทางสถิติ เช่น การถดถอย (Regression) เพื่อคาดการณ์ค่าตัวแปรตามจากตัวแปรอิสระ เช่น การใช้จำนวนชั่วโมงการอ่านหนังสือทำนายผลสอบ
- การขยายผล (Generalization) ทำให้สามารถนำผลจากกลุ่มตัวอย่างไปสรุปเชิงสถิติสำหรับประชากรได้ (Creswell & Creswell, 2018; Field, 2018)

2) ประเภทของการวิเคราะห์เชิงอนุมานที่ใช้บ่อย

(1) การทดสอบค่าเฉลี่ย (Mean Comparison)

- t-test

- *One-sample t-test* เปรียบเทียบค่าเฉลี่ยของกลุ่มเดียวกับค่ามาตรฐาน เช่น คะแนนเฉลี่ยสูงกว่าเกณฑ์ขั้นต่ำหรือไม่

ตัวอย่าง One-sample t-test

โจทย์วิจัย ต้องการตรวจสอบว่า คะแนนเฉลี่ยการสอบของนักศึกษา 30 คน สูงกว่าเกณฑ์ขั้นต่ำ 50 คะแนน หรือไม่

ข้อมูลตัวอย่าง คะแนนสอบ (N = 30)

58, 62, 55, 49, 61, 53, 59, 60, 52, 63,
54, 57, 56, 51, 65, 59, 60, 62, 55, 58,
57, 61, 54, 63, 52, 56, 60, 59, 62, 58

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู Analyze → Compare Means → One-Sample T Test
2. เลือกตัวแปร Score
3. ใส่ค่า Test Value = 50
4. กด OK → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.6 One-Sample Statistics

N	Mean	Std. Deviation	Std. Error Mean
30	57.63	3.90	0.71

ตารางที่ 8.7 One-Sample Test

Test Value = 50	t	df	Sig. (2-tailed)	Mean Difference	95% CI Lower	Upper
Score	10.75	29	.000	7.63	6.16	9.10

การตีความผลลัพธ์

- ค่าเฉลี่ยของนักศึกษา = 57.63 (SD = 3.90)
- เกณฑ์ขั้นต่ำที่กำหนด = 50
- ผลการทดสอบ $t(29) = 10.75, p < .001$
- ช่วงความเชื่อมั่น 95% ของผลต่าง = [6.16, 9.10]

สรุป คะแนนเฉลี่ยของนักศึกษากลุ่มนี้ สูงกว่าเกณฑ์ขั้นต่ำ 50 อย่างมีนัยสำคัญทางสถิติ

✍ ตัวอย่างการเขียนรายงาน (APA Style)

One-sample t-test พบว่า คะแนนเฉลี่ยของนักศึกษา ($\bar{X} = 57.63$, SD = 3.90) สูงกว่าเกณฑ์ขั้นต่ำ 50 อย่างมีนัยสำคัญทางสถิติ $t(29) = 10.75, p < .001$, 95% CI [6.16, 9.10]

- *Independent Samples t-test* เปรียบเทียบค่าเฉลี่ยของสองกลุ่มอิสระ เช่น นักเรียนชาย-หญิง

ตัวอย่าง Independent Samples t-test

โจทย์วิจัย ต้องการตรวจสอบว่า คะแนนสอบของนักเรียนชายและนักเรียนหญิง แตกต่างกันอย่างมีนัยสำคัญทางสถิติหรือไม่

ข้อมูลตัวอย่าง

- นักเรียนชาย ($n = 30$)
58, 62, 55, 49, 61, 53, 59, 60, 52, 63,
54, 57, 56, 51, 65, 59, 60, 62, 55, 58,
57, 61, 54, 63, 52, 56, 60, 59, 62, 58
- นักเรียนหญิง ($n = 30$)
64, 67, 62, 69, 70, 66, 68, 65, 63, 72,
71, 68, 67, 64, 70, 65, 66, 68, 71, 65,

67, 72, 63, 70, 69, 68, 66, 71, 70, 67

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู **Analyze** → **Compare Means** → **Independent-Samples T Test**
2. ใส่ตัวแปร **Score** เป็น *Test Variable*
3. ใส่ตัวแปร **Gender (1=ชาย, 2=หญิง)** เป็น *Grouping Variable*
4. กด **Define Groups** (1 = Male, 2 = Female)
5. คลิก **OK** → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.8 Group Statistics

Gender	N	Mean	Std. Deviation	Std. Error Mean
Male	30	57.63	3.90	0.71
Female	30	67.00	2.77	0.51

ตารางที่ 8.9 Independent Samples Test

Levene's Test F	Sig.	t	df	Sig. (2-tailed)	Mean Diff	Std. Error Diff	95% CI Lower	Upper
2.45	.124	-11.31	58	.000	-9.37	0.83	-11.04	-7.70

การตีความผลลัพธ์

- นักเรียนหญิงมีค่าเฉลี่ยคะแนนสูงกว่า ($\bar{X} = 67.00$, $SD = 2.77$)
- นักเรียนชายมีค่าเฉลี่ยต่ำกว่า ($\bar{X} = 57.63$, $SD = 3.90$)
- ผลการทดสอบ: $t(58) = -11.31$, $p < .001$
- ผลต่างเฉลี่ย = -9.37 คะแนน
- 95% CI = $[-11.04, -7.70]$

สรุป คะแนนสอบเฉลี่ยของนักเรียนหญิง สูงกว่านักเรียนชายอย่างมีนัยสำคัญทางสถิติ

✍ ตัวอย่างการเขียนรายงาน (APA Style)

Independent-samples t-test พบว่า คะแนนสอบเฉลี่ยของนักเรียนหญิง ($\bar{X} = 67.00$, $SD = 2.77$) สูงกว่านักเรียนชาย ($\bar{X} = 57.63$, $SD = 3.90$) อย่างมีนัยสำคัญทางสถิติ $t(58) = -11.31$, $p < .001$, 95% CI $[-11.04, -7.70]$, Cohen's $d = 2.89$

- *Paired Samples t-test* เปรียบเทียบค่าเฉลี่ยก่อน-หลังของกลุ่มเดียวกัน เช่น คะแนนก่อน-หลังเรียน

ตัวอย่าง Paired Samples t-test

โจทย์วิจัย ต้องการตรวจสอบว่า คะแนนสอบก่อนเรียน (Pre-test) และหลังเรียน (Post-test) ของนักเรียนกลุ่มเดียวกัน (N = 30) แตกต่างกันหรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.10 คะแนนสอบของนักเรียน

นักเรียน	Pre-test	Post-test
1	52	60
2	55	63
3	50	58
4	48	55
5	60	68
...	...	...
30	54	62

(ค่าเฉลี่ย Pre = 54.20, Post = 62.10)

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู Analyze → Compare Means → Paired-Samples T Test
2. เลือกตัวแปร Pre-test และ Post-test จับคู่ในช่อง Paired Variables
3. คลิก OK → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.11 Paired Samples Statistics

	Mean	N	Std. Deviation	Std. Error Mean
Pre	54.20	30	4.12	0.75
Post	62.10	30	3.95	0.72

ตารางที่ 8.12 Paired Samples Test

	Mean Difference	t	df	Sig. (2-tailed)	95% CI Lower	Upper
Pre-Post	-7.90	-12.35	29	.000	-9.20	-6.60

การตีความผลลัพธ์

- คะแนนหลังเรียน ($\bar{X}$ = 62.10, SD = 3.95) สูงกว่าคะแนนก่อนเรียน ($\bar{X}$ = 54.20, SD = 4.12)

- ผลการทดสอบ: $t(29) = -12.35$, $p < .001$

- ผลต่างเฉลี่ย = -7.90 คะแนน

- ช่วงความเชื่อมั่น 95% ของผลต่าง = [-9.20, -6.60]

สรุป คะแนนเฉลี่ยหลังเรียน สูงวกาก่อนเรียนอย่างมีนัยสำคัญทาง

สถิติ

✎ ตัวอย่างการเขียนรายงาน (APA Style)

Paired-samples t-test พบว่าคะแนนหลังเรียน ($\bar{X}$ = 62.10, SD = 3.95) สูงกว่าคะแนนก่อนเรียน ($\bar{X}$ = 54.20, SD = 4.12) อย่างมีนัยสำคัญทางสถิติ, $t(29) = -12.35$, $p < .001$, 95% CI [-9.20, -6.60], Cohen's $d = 2.25$

- ANOVA (Analysis of Variance)

- *One-way ANOVA* เปรียบเทียบค่าเฉลี่ยมากกว่า 2 กลุ่มอิสระ เช่น ระดับการศึกษาต่ำ-กลาง-สูง

ตัวอย่าง One-way ANOVA

โจทย์วิจัย ต้องการตรวจสอบว่า ระดับการศึกษา (ต่ำ-กลาง-สูง)

มีผลต่อ คะแนนสอบ ของนักศึกษาแตกต่างกันหรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.13 ระดับการศึกษา (ต่ำ-กลาง-สูง)

กลุ่มการศึกษา	N	คะแนนสอบเฉลี่ย (Mean)	SD
ต่ำ (High school)	20	55.2	4.1
กลาง (Bachelor)	20	60.3	3.8
สูง (Graduate)	20	63.5	3.2

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู **Analyze** → **Compare Means** → **One-Way ANOVA**
2. ใส่ตัวแปร **Score** เป็น *Dependent Variable*
3. ใส่ตัวแปร **Education Level (ต่ำ/กลาง/สูง)** เป็น *Factor*
4. กด **Post Hoc** (เลือก **Tukey**) หากต้องการดูว่ากลุ่มไหนต่างกัน
5. คลิก **OK** → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.14 ANOVA

Source	SS	df	MS	F	Sig.
Between Groups	780.27	2	390.14	10.85	.000
Within Groups	2016.50	57	35.39		
Total	2796.77	59			

ตารางที่ 8.15 Post Hoc Tests (Tukey HSD)

(I) Education	(J) Education	Mean Diff (I-J)	Sig.
Low	Middle	-5.10	.012
Low	High	-8.30	.000
Middle	High	-3.20	.045

การตีความผลลัพธ์

- ค่าทดสอบ ANOVA: $F(2,57) = 10.85, p < .001$ → มีความแตกต่างอย่างมีนัยสำคัญทางสถิติระหว่างกลุ่ม
- การเปรียบเทียบ Post Hoc (Tukey) พบว่า:
 - กลุ่มการศึกษาสูงมีคะแนนเฉลี่ย สูงกว่ากลุ่มต่ำ อย่างมีนัยสำคัญ ($p < .001$)
 - กลุ่มการศึกษากลางมีคะแนนสูงกว่ากลุ่มต่ำอย่างมีนัยสำคัญ ($p = .012$)
 - กลุ่มสูงสูงกว่ากลุ่มกลางเล็กน้อย แต่ยังมีนัยสำคัญ ($p = .045$)

สรุป ระดับการศึกษาที่มีผลต่อคะแนนสอบ โดยระดับการศึกษาที่สูงกว่าจะมีผลสัมฤทธิ์ทางการเรียนเฉลี่ยสูงกว่า

✍ ตัวอย่างการเขียนรายงาน (APA Style)

One-way ANOVA พบว่าคะแนนสอบแตกต่างกันตามระดับการศึกษา $F(2, 57) = 10.85, p < .001$. การทดสอบ Post Hoc (Tukey) ชี้ว่า นักศึกษาระดับปริญญาโท/เอกมีคะแนนสูงกว่ากลุ่มปริญญาตรีและมัธยมอย่างมีนัยสำคัญ

- *Two-way ANOVA* ใช้เมื่อมีตัวแปรอิสระ 2 ตัว และต้องการดูผลของแต่ละตัวรวมถึงปฏิสัมพันธ์ (Interaction Effect)

ตัวอย่าง Two-way ANOVA

โจทยวิจัย ผู้วิจัยต้องการตรวจสอบว่า เพศ (ชาย-หญิง) และ วิธีการสอน (ดั้งเดิม-ใหม่) มีผลต่อ คะแนนสอบหลังเรียน (Post-test) หรือไม่ และต้องการดูว่าเพศกับวิธีการสอนมี ปฏิสัมพันธ์ (Interaction Effect) กันหรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.16 เพศกับวิธีการสอน

กลุ่ม	เพศ	วิธีการสอน	N	Mean	SD
1	ชาย	ดั้งเดิม	15	58.2	4.5
2	ชาย	ใหม่	15	64.1	3.9
3	หญิง	ดั้งเดิม	15	60.5	4.2
4	หญิง	ใหม่	15	68.0	3.6

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่ Analyze → General Linear Model → Univariate
2. ใส่ Score_Post เป็น *Dependent Variable*
3. ใส่ Gender และ Method เป็น *Fixed Factors*
4. คลิก Model → Full Factorial เพื่อรวมปฏิสัมพันธ์ (Gender × Method)
5. คลิก Plots หากต้องการดู Interaction graph
6. คลิก OK → ได้ผลลัพธ์ Output

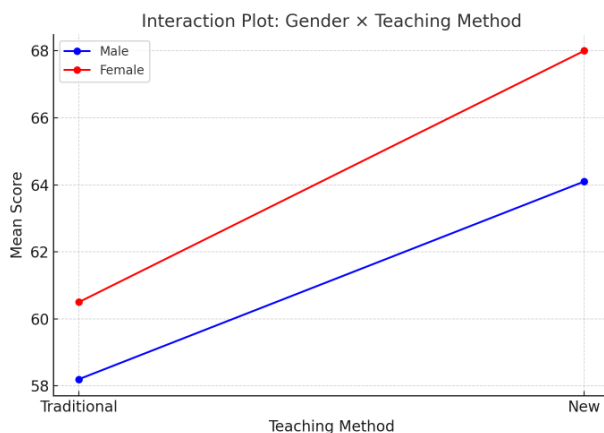
ตัวอย่าง Output (SPSS)

ตารางที่ 8.17 Tests of Between-Subjects Effects

Source	SS	df	MS	F	Sig.
Gender	210.25	1	210.25	12.40	.001
Method	512.64	1	512.64	30.22	.000
Gender × Method	36.50	1	36.50	2.15	.148
Error	1018.43	56	18.18		
Total	24765.0	60			

แผนภาพ Interaction Plot

- เส้นของกลุ่มเพศชายและหญิง ขนานกันเกือบสนิท → แสดงว่า ไม่มีปฏิสัมพันธ์ชัดเจน
- แต่ระดับเส้นต่างกัน (ชายต่ำกว่า หญิงสูงกว่า) → แสดงว่า เพศมีผลหลัก (Main Effect)
- และเส้นทั้งคู่สูงขึ้นเมื่อใช้วิธีการสอนใหม่ → แสดงว่า วิธีการสอนมีผลหลัก (Main Effect)



แผนภาพที่ 8.8 Interaction Plot

การตีความผลลัพธ์

- Gender มีผลต่อคะแนนสอบ $F(1,56) = 12.40$, $p = .001$ → หญิงได้คะแนนสูงกว่าชาย
- Method มีผลต่อคะแนนสอบ $F(1,56) = 30.22$, $p < .001$ → การสอนแบบใหม่ดีกว่าดั้งเดิม

- Interaction (Gender $\times$ Method) ไม่พบปฏิสัมพันธ์ $F(1,56) = 2.15, p = .148 \rightarrow$ ผลของวิธีสอนไม่ขึ้นอยู่กับเพศ

สรุป ทั้งเพศและวิธีการสอนมีอิทธิพลต่อคะแนนสอบ แต่ไม่มีผลร่วมกันอย่างมีนัยสำคัญ

๔ ตัวอย่างการเขียนรายงาน (APA Style)

Two-way ANOVA พบว่า เพศมีผลหลักต่อคะแนนสอบ $F(1,56) = 12.40, p = .001, \eta^2 = .18$. วิธีการสอนก็มีผลหลักอย่างมีนัยสำคัญ $F(1,56) = 30.22, p < .001, \eta^2 = .35$. อย่างไรก็ตาม ไม่พบผลปฏิสัมพันธ์ระหว่างเพศกับวิธีการสอน $F(1,56) = 2.15, p = .148$

(2) การวิเคราะห์ความสัมพันธ์ (Correlation)

- Pearson's r ใช้สำหรับข้อมูลช่วง/อัตราส่วน (Interval/Ratio) และเป็นการแจกแจงปกติ

ตัวอย่าง Pearson's r

โจทยวิจัย ผู้วิจัยต้องการตรวจสอบว่า จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์ มีความสัมพันธ์กับ คะแนนสอบปลายภาค ของนักศึกษา หรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.18 จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์และคะแนนสอบปลายภาคของนักศึกษา

นักศึกษา	ชั่วโมงอ่านหนังสือ (Hours)	คะแนนสอบ (Score)
1	5	52
2	8	60
3	6	55
4	10	65
5	12	70
6	7	58
7	15	75
8	9	63
9	11	68
10	13	72

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู Analyze → Correlate → Bivariate
2. เลือกตัวแปร Hours และ Score
3. ดึงเลือก Pearson และ Two-tailed
4. คลิก OK → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.19 Tests of Between-Subjects Effects

Variables	Hours	Score
Hours	1	.92**
Score	.92**	1

หมายเหตุ: $p < .001$

การตีความผลลัพธ์

- ค่า $r = .92$, $p < .001$
- แสดงว่าจำนวนชั่วโมงอ่านหนังสือ มีความสัมพันธ์เชิงบวกสูงมากกับคะแนนสอบ

สรุป นักศึกษาที่อ่านหนังสือมาก มักจะได้คะแนนสอบสูงขึ้น

✍ ตัวอย่างการเขียนรายงาน (APA Style)

การวิเคราะห์ค่าสหสัมพันธ์แบบเพียร์สัน (Pearson's r) พบว่าชั่วโมงอ่านหนังสือต่อสัปดาห์มีความสัมพันธ์เชิงบวกสูงกับคะแนนสอบ, $r(8) = .92$, $p < .001$

สรุป Pearson's r ใช้หาความสัมพันธ์เชิงเส้นตรงของตัวแปร ช่วง/อัตราส่วน (Interval/Ratio) ที่มีการแจกแจงปกติ

- ค่า r อยู่ระหว่าง -1 ถึง +1
- $r > 0$ → ความสัมพันธ์เชิงบวก
- $r < 0$ → ความสัมพันธ์เชิงลบ
- $|r|$ ยิ่งใกล้ 1 → ความสัมพันธ์ยิ่งแน่นแฟ้น
- Spearman's ρ ใช้สำหรับข้อมูลลำดับ (Ordinal) หรือข้อมูลที่ไม่เป็นปกติ

ตัวอย่าง Spearman's ρ

โจทย์วิจัย ผู้วิจัยต้องการตรวจสอบว่า ระดับความพึงพอใจต่อการเรียน (Likert 1–5) มีความสัมพันธ์กับ อันดับผลสัมฤทธิ์ทางการเรียน (1 = อันดับสูงสุด, 10 = อันดับต่ำสุด) หรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.20 ระดับความพึงพอใจต่อการเรียน

นักศึกษา	ความพึงพอใจ (Satisfaction)	อันดับผลสัมฤทธิ์ (Rank)
1	5	1
2	4	3
3	3	7
4	4	5
5	2	9
6	5	2
7	3	6
8	1	10
9	4	4
10	2	8

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู Analyze → Correlate → Bivariate
2. เลือกตัวแปร Satisfaction และ Rank
3. ตี๊กเลือก Spearman (ไม่เลือก Pearson)
4. คลิก OK → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.21 Tests of Between-Subjects Effects

Variables	Satisfaction	Rank
Satisfaction	1	-.88**
Rank	-.88**	1

หมายเหตุ: $p < .001$

การตีความผลลัพธ์

- ค่า $\rho = -0.88$, $p < .001$
- แสดงว่าความพึงพอใจมี **ความสัมพันธ์เชิงลบ**สูงมาก กับลำดับผลสัมฤทธิ์
- สรุป ยิ่งนักศึกษามีความพึงพอใจต่อการเรียนมาก $\rightarrow$ ยิ่งมีอันดับผลสัมฤทธิ์ที่ต่ำกว่า (อันดับตัวเลขน้อยกว่า)

✍ ตัวอย่างการเขียนรายงาน (APA Style)

การวิเคราะห์สหสัมพันธ์แบบสเปียร์แมน (Spearman's rho) พบว่าความพึงพอใจในการเรียนสัมพันธ์เชิงลบกับอันดับผลสัมฤทธิ์, $\rho(8) = -0.88$, $p < .001$

สรุป Spearman's rho เหมาะใช้กับ

- ข้อมูล ลำดับ (Ordinal data) เช่น Likert scale
- ข้อมูลที่ **ไม่เป็นปกติ** (Non-normal distribution)
- หรือเมื่อมี **outlier** ที่อาจกระทบการวิเคราะห์แบบ Pearson

(3) การวิเคราะห์การถดถอย (Regression Analysis)

- Simple Regression ใช้ตัวแปรอิสระ 1 ตัวทำนายตัวแปรตาม

ตัวอย่าง Simple Regression

โจทยวิจัย ผู้วิจัยต้องการตรวจสอบว่า จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์ (Hours) สามารถใช้ทำนาย คะแนนสอบ (Score) ได้หรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.22 จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์

นักศึกษา	ชั่วโมงอ่าน (Hours)	คะแนนสอบ (Score)
1	5	52
2	8	60
3	6	55
4	10	65
5	12	70
6	7	58
7	15	75
8	9	63

นักศึกษา	ชั่วโมงอ่าน (Hours)	คะแนนสอบ (Score)
9	11	68
10	13	72

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู **Analyze** → **Regression** → **Linear**
2. ใส่ตัวแปร **Score** เป็น *Dependent*
3. ใส่ตัวแปร **Hours** เป็น *Independent*
4. คลิก **OK** → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.23 Model Summary

R	R ²	Adjusted R ²	Std. Error
.92	.85	.84	2.24

ตารางที่ 8.24 ANOVA Table

Source	SS	df	MS	F	Sig.
Regression	512.6	1	512.6	102.5	.000
Residual	91.2	8	11.4		
Total	603.8	9			

ตารางที่ 8.25 Coefficients Table

Predictor	B	Std. Error	Beta	t	Sig.
Constant	40.12	2.11	–	18.99	.000
Hours	2.50	0.25	.92	10.12	.000

การตีความผลลัพธ์

- ค่า $R^2 = .85$ → โมเดลสามารถอธิบายความแปรปรวนของคะแนนสอบได้ 85%
- ค่าความชัน (B) ของ Hours = 2.50, $p < .001$ → ทุก ๆ การอ่านเพิ่มขึ้น 1 ชั่วโมง คะแนนสอบจะเพิ่มขึ้นเฉลี่ย 2.5 คะแนน
- สมการถดถอยที่ได้

$$\text{Score} = 40.12 + 2.50 (\text{Hours})$$

๔ ตัวอย่างการเขียนรายงาน (APA Style)

Simple linear regression พบว่า จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์ สามารถทำนายคะแนนสอบได้อย่างมีนัยสำคัญทางสถิติ $F(1,8) = 102.5, p < .001, R^2 = .85$ สมการถดถอยคือ $\text{Score} = 40.12 + 2.50(\text{Hours})$

สรุป Simple Regression เป็นการสร้างสมการเชิงเส้นเพื่อทำนายค่าตัวแปรตามจากตัวแปรอิสระเพียงตัวเดียว

- ต้องตรวจสอบสมมติฐาน (linear relationship, normality ของ residuals, homoscedasticity)
- ผลลัพธ์สามารถนำไปใช้ในการพยากรณ์และตีความเชิงวิจัยได้
- Multiple Regression ใช้หลายตัวแปรอิสระเพื่ออธิบายความแปรปรวนของตัวแปรตาม

ตัวอย่าง Multiple Regression

โจทยวิจัย ผู้วิจัยต้องการตรวจสอบว่า จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์ (Hours) และคะแนนก่อนเรียน (Pre-test) สามารถร่วมกันทำนาย คะแนนสอบหลังเรียน (Post-test) ได้หรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.26 จำนวนชั่วโมงอ่านหนังสือต่อสัปดาห์

นักศึกษา	Hours	Score_Pre	Score_Post
1	5	50	55
2	8	54	63
3	6	52	58
4	10	57	66
5	12	59	71
6	7	53	60
7	15	62	77
8	9	56	65
9	11	58	70
10	13	60	73

ขั้นตอนการวิเคราะห์ (SPSS)

1. ไปที่เมนู **Analyze** → **Regression** → **Linear**
2. ใส่ตัวแปร **Score_Post** เป็น *Dependent*
3. ใส่ตัวแปร **Hours, Score_Pre** เป็น *Independent(s)*
4. คลิก **Statistics** → เลือก *Estimates, R² change, Collinearity diagnostics*
5. คลิก **OK** → ได้ผลลัพธ์ Output

ตัวอย่าง Output (SPSS)

ตารางที่ 8.27 Model Summary

R	R ²	Adjusted R ²	Std. Error
.92	.85	.83	2.10

ตารางที่ 8.28 ANOVA Table

Source	SS	df	MS	F	Sig.
Regression	502.6	2	251.3	56.9	.000
Residual	88.2	7	12.6		
Total	590.8	9			

ตารางที่ 8.29 Coefficients Table

Predictor	B	Std. Error	Beta	t	Sig.
Constant	20.15	3.50	–	5.76	.001
Hours	1.80	0.40	.62	4.50	.003
Score_Pre	0.60	0.15	.55	4.00	.005

การตีความผลลัพธ์

- ค่า $R^2 = .85$ → โมเดลสามารถอธิบายความแปรปรวนของคะแนนหลังเรียนได้ **85%**
- ตัวแปร Hours ($B = 1.80$, $p = .003$) และ Score_Pre ($B = 0.60$, $p = .005$) มีอิทธิพลต่อ Score_Post อย่างมีนัยสำคัญ
- สมการถดถอยที่ได้

$$\text{Score_Post} = 20.15 + 1.80 (\text{Hours}) + 0.60 (\text{Score_Pre})$$

๔ ตัวอย่างการเขียนรายงาน (APA Style)

การวิเคราะห์ Multiple Regression พบว่าจำนวนชั่วโมงอ่านหนังสือและคะแนนก่อนเรียนสามารถร่วมกันทำนายคะแนนหลังเรียนได้อย่างมีนัยสำคัญทางสถิติ $F(2,7) = 56.9, p < .001, R^2 = .85$ สมการถดถอยที่ได้คือ $\text{Score_Post} = 20.15 + 1.80(\text{Hours}) + 0.60(\text{Score_Pre})$

สรุป **Multiple Regression** เหมาะเมื่อผู้วิจัยต้องการใช้ตัวแปรอิสระหลายตัวร่วมกันในการทำนายตัวแปรตาม

- ช่วยอธิบายความแปรปรวนได้มากกว่า Simple Regression
- ต้องตรวจสอบ **Multicollinearity** (เช่น VIF) และสมมติฐานพื้นฐานอื่น ๆ (Linearity, Normality, Homoscedasticity)

(4) การทดสอบไคสแควร์ (Chi-square Test)

การทดสอบ Chi-square ใช้สำหรับข้อมูลเชิงจัดประเภท (Nominal/Ordinal) เพื่อตรวจสอบว่าตัวแปร 2 ตัวมีความสัมพันธ์กันหรือไม่ โดยเปรียบเทียบความถี่ที่สังเกตได้ (Observed frequency) กับความถี่ที่คาดหวัง (Expected frequency) หากค่าที่แตกต่างกันมากจนเกินไป จะได้สถิติไคสแควร์ที่มีนัยสำคัญ

ตัวอย่างโจทย์ ต้องการตรวจสอบว่า เพศ (ชาย-หญิง) มีความสัมพันธ์กับการเลือกวิชาเลือก (A, B, C) หรือไม่

ข้อมูลตัวอย่าง

ตารางที่ 8.30 Crosstab (Observed frequencies)

	วิชา A	วิชา B	วิชา C	รวม
ชาย (Male)	10	20	15	45
หญิง (Female)	25	15	15	55
รวม	35	35	30	100

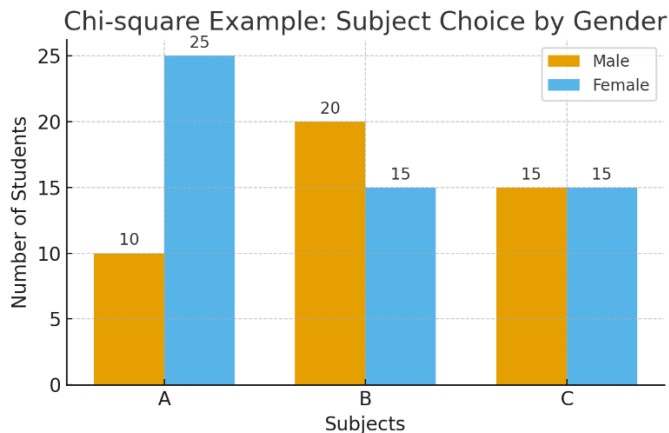
ตารางที่ 8.31 Chi-Square Tests

Test	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	12.56	2	.002

แผนภาพตัวอย่างผลการทดสอบไคสแควร์ (Chi-square Test) ที่แสดงการเลือกวิชาเรียน (Subjects A, B, C) จำแนกตามเพศ (ชาย-หญิง)

- จะเห็นว่าเพศหญิงเลือกวิชา A มากกว่าชาย

- เพศชายเลือกวิชา B มากกว่าเพศหญิง
- ค่าการทดสอบ Chi-square ($\chi^2 = 12.56$, $df = 2$, $p = .002$) บ่งชี้ว่า เพศมีความสัมพันธ์กับการเลือกวิชาเรียนอย่างมีนัยสำคัญทางสถิติ



แผนภาพที่ 8.9 ตัวอย่างผลการทดสอบไคสแควร์ (Chi-square Test)

การตีความผลลัพธ์

- ค่า Chi-square = 12.56, $df = 2$, $p = .002$
- ค่า $p < .05 \rightarrow$ สรุปว่าเพศกับการเลือกวิชาเรียน มีความสัมพันธ์กันอย่างมีนัยสำคัญทางสถิติ
- กล่าวคือ เพศมีผลต่อแนวโน้มการเลือกวิชาเรียน

✍ ตัวอย่างการเขียนรายงาน (APA Style)

การทดสอบไคสแควร์ (Chi-square test of independence) พบว่า เพศมีความสัมพันธ์กับการเลือกวิชาเรียนอย่างมีนัยสำคัญทางสถิติ, $\chi^2(2, N = 100) = 12.56$, $p = .002$

สรุป Chi-square test เหมาะสำหรับข้อมูล Nominal/Ordinal

- ใช้ทดสอบความสัมพันธ์ของ สองตัวแปรเชิงจัดประเภท
- ต้องระวัง cell ที่มี expected frequency < 5 ไม่ควรมากกว่า 20% ของเซลล์ทั้งหมด (ควรใช้ Fisher's Exact Test แทน)

สรุปได้ว่า สถิติอนุมานทำให้นักวิจัยก้าวข้ามจากข้อมูลเชิงพรรณนาไปสู่การอธิบาย สรุป และทำนาย ด้วยหลักการทางคณิตศาสตร์และความน่าจะเป็น การเลือกใช้สถิติที่ถูกต้อง ร่วมกับการตีความอย่างรอบคอบและเชื่อมโยงกับทฤษฎี จะทำให้นักวิจัยมีคุณค่าเชิงวิชาการ และ น่าเชื่อถือ ทั้งในการเรียนการสอนและการประยุกต์ใช้จริง

5. การแปลผล Output

การวิเคราะห์ข้อมูลทางสถิติด้วยโปรแกรม (เช่น SPSS, JASP, R) จะได้ผลลัพธ์ในรูปแบบตารางและกราฟ แต่การนำผลเหล่านั้นไปใช้จำเป็นต้องตีความเชิงความหมาย ไม่ใช่เพียงการรายงานตัวเลขดิบ

1) หลักการสำคัญของการแปลผล

(1) ค่าความน่าจะเป็น (p-value)

- $p < .05 \rightarrow$ แสดงว่ามี **นัยสำคัญทางสถิติ** (reject H_0)
- $p \geq .05 \rightarrow$ แสดงว่า **ไม่มีนัยสำคัญ** (fail to reject H_0)
- แต่ p-value ไม่บอกขนาดของผล (Effect Size) และไม่ใช้ความจริงของสมมติฐาน

(2) ค่าตัวสถิติ (Test Statistic)

- t (t-test), F (ANOVA/Regression), χ^2 (Chi-square), r (Correlation) $\rightarrow$ ใช้ยืนยันผล
- ต้องรายงานพร้อมค่า df (degrees of freedom) เช่น $t(103) = 2.45$, $p < .05$

(3) ขนาดอิทธิพล (Effect Size)

- ใช้เสริม p-value เพื่อบอกว่าผลลัพธ์มีความสำคัญในเชิงปฏิบัติหรือไม่
- ตัวชี้วัดที่ใช้บ่อย เช่น
 - Cohen's d (สำหรับ t-test): 0.2 = เล็ก, 0.5 = กลาง, 0.8 = ใหญ่
 - η^2 หรือ partial η^2 (สำหรับ ANOVA): 0.01 = เล็ก, 0.06 = กลาง, 0.14 = ใหญ่
 - R^2 (สำหรับ Regression): สัดส่วนความแปรปรวนที่โมเดลอธิบายได้

(4) ช่วงความเชื่อมั่น (Confidence Interval, CI)

- แสดงช่วงค่าที่คาดว่า “ค่าจริงของประชากร” จะอยู่ภายใน เช่น 95% CI [38.2, 41.5]

(5) การเชื่อมโยงกับคำถามวิจัย

- ไม่ควรหยุดแค่บอกว่า “มีนัยสำคัญ” แต่ควรตอบว่า ผลนี้บอกอะไรต่อสมมติฐาน

2) ขั้นตอนการอ่าน Output

1. ดูตารางสรุป (Descriptives) $\rightarrow \bar{X}$, SD, จำนวนตัวอย่าง
2. ตรวจสอบ Assumptions $\rightarrow$ เช่น Levene's test สำหรับ t-test

3. อ่านค่าตัวสถิติหลัก → เช่น t , F , r , R^2
4. ดูค่า p -value → ตัดสินใจยอมรับ/ปฏิเสธสมมติฐาน
5. เสริมด้วย Effect Size และ CI
6. ตีความเชิงความหมาย → แปลว่ากลุ่มไหนดีกว่า ตัวแปรสัมพันธ์กันอย่างไร

3) ตัวอย่างการแปลผล

(1) Independent Samples t-test

Output (SPSS):

- กลุ่ม A: $\bar{X} = 38.63$, $SD = 6.31$
- กลุ่ม B: $\bar{X} = 41.87$, $SD = 9.36$
- $t(103.6) = -2.22$, $p = .028$, Cohen's $d = 0.39$

การตีความ นักเรียนกลุ่ม B มีคะแนนสูงกว่ากลุ่ม A อย่างมีนัยสำคัญทางสถิติ ($p < .05$) และมีขนาดอิทธิพลระดับเล็กถึงกลาง ($d = 0.39$) แสดงว่าวิธีการสอนที่ใช้กับกลุ่ม B อาจมีผลต่อการพัฒนาคะแนน

(2) Pearson Correlation

Output (JASP/R):

- $r = .53$, $p < .001$

การตีความ จำนวนชั่วโมงอ่านหนังสือมีความสัมพันธ์เชิงบวกปานกลางกับคะแนนหลังเรียน ซึ่งหมายความว่า นักเรียนที่อ่านหนังสือมากกว่ามีแนวโน้มที่จะได้คะแนนสูงขึ้น

(3) Regression

Output:

- Predictors: Hours ($\beta = .47$, $p < .001$), Score_Pre ($\beta = .31$, $p = .001$)
- $R^2 = .42$

การตีความ ชั่วโมงอ่านหนังสือและคะแนนก่อนเรียนเป็นตัวทำนายคะแนนหลังเรียนที่มีนัยสำคัญ โดยโมเดลสามารถอธิบายความแปรปรวนของคะแนนหลังเรียนได้ถึง 42%

(4) ANOVA

Output:

- $F(2, 116) = 1.81$, $p = .17$, $\eta^2 = .03$

การตีความ ค่าเฉลี่ยผลสัมฤทธิ์ของนักเรียน 3 กลุ่มการศึกษามีแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($p > .05$) และมีขนาดอิทธิพลเล็ก ($\eta^2 = .03$)

สรุปได้ว่า การแปลผล Output ควร 1. รายงานค่าตัวเลขหลักครบถ้วน (t , F , χ^2 , r , p , df , *effect size*, *CI*) 2. อธิบายความหมายในเชิงวิชาการ ไม่ใช่แค่รายงานค่า p และ 3. เชื่อมโยงผลลัพธ์กลับไปยัง สมมติฐานและคำถามวิจัย

6. การนำเสนอผลลัพธ์ในรูปแบบข้อความ ตาราง และกราฟ

เมื่อวิเคราะห์ข้อมูลเสร็จ ขั้นตอนต่อมาคือ การสื่อสารผลลัพธ์เป็นขั้นตอนสำคัญที่ทำให้ผู้อ่าน เข้าใจง่าย ถูกต้อง และเป็นไปตามมาตรฐานวิชาการ การนำเสนอที่ดี ควรผสมผสานทั้ง ข้อความ (Text), ตาราง (Tables) และ กราฟ/แผนภาพ (Figures) เพื่อช่วยอธิบายผลให้สมบูรณ์

1) การนำเสนอผลลัพธ์ในรูปแบบข้อความ (Text Presentation)

- ควรอธิบายผลอย่างกระชับ เชิงความหมาย ไม่ใช่การคัดลอกตัวเลขจาก Output
- ใช้รูปแบบ APA Style เช่น ระบุค่าเฉลี่ย ($\bar{X}$), ส่วนเบี่ยงเบนมาตรฐาน (SD), ค่า test statistic, df , p -value และ effect size

ตัวอย่าง ค่าเฉลี่ยคะแนนหลังเรียนของกลุ่ม B ($\bar{X} = 41.87$, $SD = 9.36$) สูงกว่ากลุ่ม A ($\bar{X} = 38.63$, $SD = 6.31$) อย่างมีนัยสำคัญทางสถิติ $t(103.6) = -2.22$, $p = .028$, Cohen's $d = 0.39$

2) การนำเสนอผลลัพธ์ในรูปแบบตาราง (Tables)

- ตารางช่วยสรุปข้อมูลจำนวนมากให้กระชับ อ่านง่าย
- ควรใส่ เลขตาราง + ชื่อเรื่อง (Caption) ตามมาตรฐาน เช่น “ตารางที่ 1 ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของคะแนนก่อน-หลังเรียน”
- ใช้รูปแบบที่เรียบง่าย ไม่ซับซ้อน

ตัวอย่าง

ตารางที่ 8.32 ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของคะแนนก่อน-หลังเรียน

Group	N	Mean (M)	SD	t(df)	p
A	59	38.63	6.31		
B	60	41.87	9.36	-2.22 (103.6)	.028

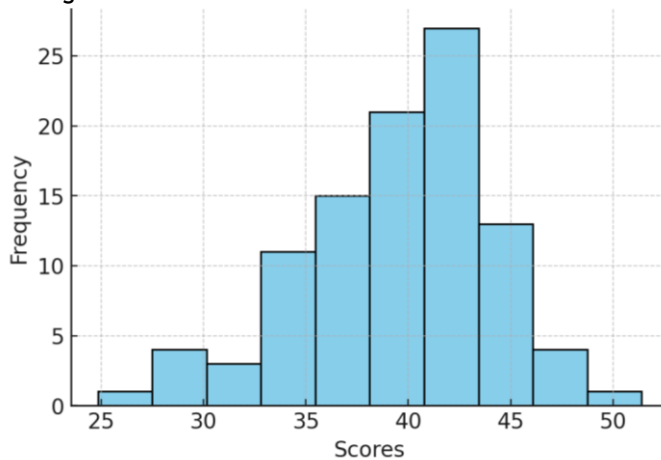
การตีความ ผลการทดสอบชี้ให้เห็นว่ากลุ่ม B มีค่าเฉลี่ยสูงกว่ากลุ่ม A อย่างมีนัยสำคัญ

3) การนำเสนอผลลัพธ์ในรูปกราฟ (Figures)

- กราฟทำให้เห็น แนวโน้ม ความแตกต่าง หรือความสัมพันธ์ ได้ชัดเจน
- ประเภทกราฟที่นิยมใช้ในงานวิชาการ ได้แก่
 - Histogram → การกระจายของคะแนน
 - Bar Chart → ค่าเฉลี่ยของแต่ละกลุ่ม
 - Boxplot → การเปรียบเทียบการกระจายและค่าเฉลี่ยของหลายกลุ่ม
 - Scatter Plot → ความสัมพันธ์ระหว่างสองตัวแปร
 - Line Chart → การเปลี่ยนแปลงตามเวลา

ตัวอย่าง

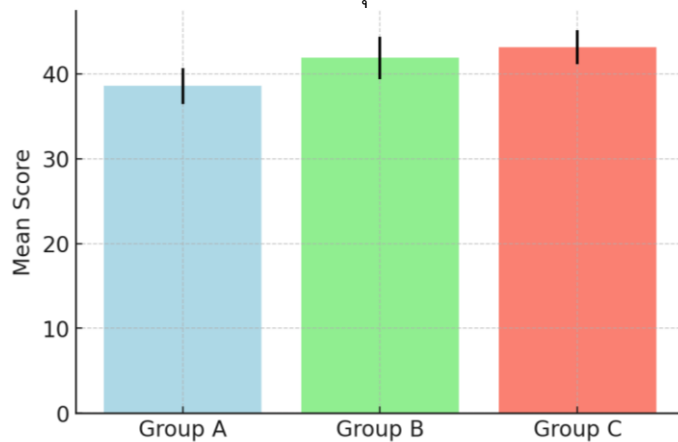
- Histogram การกระจายของคะแนนหลังเรียน



แผนภาพที่ 8.10 แผนภาพตัวอย่าง Histogram

- กราฟแท่งแสดงการกระจายของคะแนนที่นักเรียนได้หลังเรียน
- การแจกแจงใกล้เคียงกับโค้งปกติ (Normal Distribution)
- เห็นค่ากลาง (Mean) อยู่ใกล้ช่วง 40 คะแนน

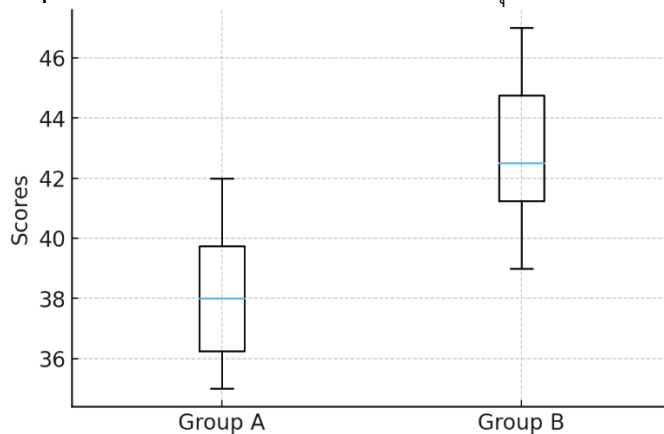
- Bar Chart ค่าเฉลี่ยของแต่ละกลุ่ม



แผนภาพที่ 8.11 แผนภาพตัวอย่าง Bar Chart

- แสดงค่าเฉลี่ยคะแนนหลังเรียนของกลุ่ม A, B และ C
- กลุ่ม C มีค่าเฉลี่ยสูงที่สุด รองลงมาคือกลุ่ม B และ A
- Error bar แสดงค่าความแปรปรวน (SD/SE) ของแต่ละกลุ่ม

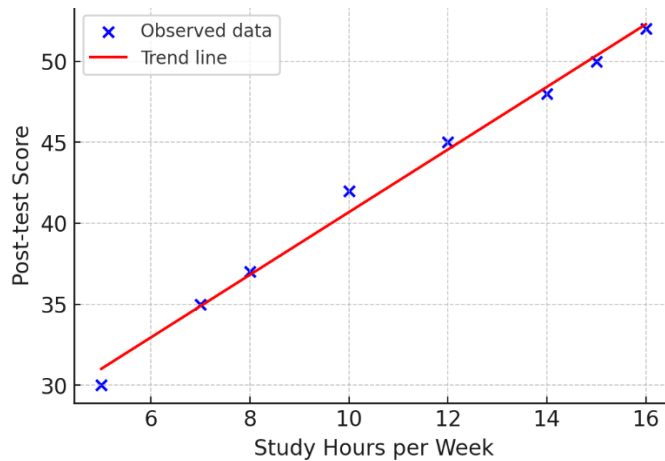
- Boxplot เปรียบเทียบคะแนนหลังเรียนของกลุ่ม A และ B



แผนภาพที่ 8.12 แผนภาพตัวอย่าง Boxplot

- แสดงว่ากลุ่ม B มีค่ามัธยฐานสูงกว่ากลุ่ม A
- การกระจายของข้อมูลกลุ่ม B กว้างกว่าเล็กน้อย

- Scatter Plot ความสัมพันธ์ระหว่างชั่วโมงอ่านหนังสือกับคะแนนหลังเรียน

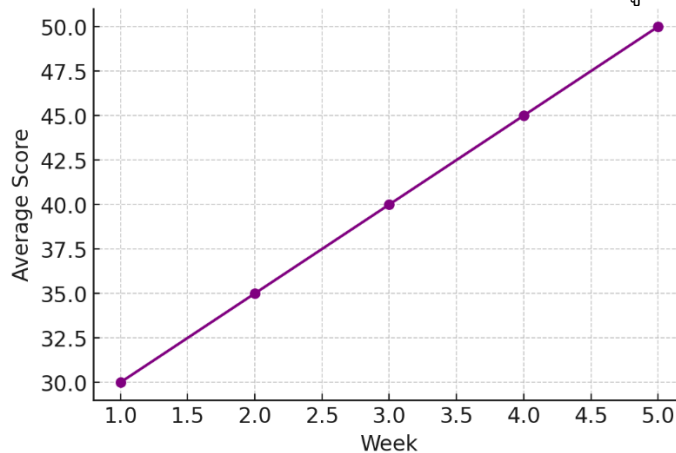


แผนภาพที่ 8.13 แผนภาพตัวอย่าง Scatter Plot

- จุดข้อมูลแสดงการกระจายสัมพันธ์ระหว่างจำนวนชั่วโมงอ่านหนังสือกับคะแนนหลังเรียน
- เส้นแนวโน้มชี้ขึ้น แสดงว่ามีความสัมพันธ์เชิงบวก ($r = .53, p < .001$)
- นักเรียนที่อ่านมากกว่ามีแนวโน้มได้คะแนนสูงกว่า

- Line Chart การเปลี่ยนแปลงของคะแนนเฉลี่ยตามเวลา

- เส้นกราฟแสดงพัฒนาการคะแนนเฉลี่ยในช่วง 5 สัปดาห์
- เห็นการเพิ่มขึ้นต่อเนื่องของผลสัมฤทธิ์การเรียนรู้
- สะท้อนให้เห็นแนวโน้มเชิงบวกจากการสอน/การเรียนรู้ต่อเนื่อง



แผนภาพที่ 8.14 แผนภาพตัวอย่าง Line Chart

4) หลักการเขียน Caption และอ้างอิง

- (1) หมายเลข + ชื่อกำกับ (Numbering & Title)
 - ตารางใช้คำว่า “ตารางที่ ...” (Table ...)
 - แผนภาพ/กราฟใช้คำว่า “แผนภาพที่ ...” หรือ “รูปที่ ...” (Figure ...)
 - ควรเรียงหมายเลขตามลำดับที่ปรากฏในบท
- (2) คำบรรยายใต้ภาพ/ตาราง (Caption/Legend)
 - เขียนสั้น กระชับ และชี้ชัดว่ากำลังแสดงอะไร
 - ถ้าเป็นตารางแนะนำให้ใส่ “ตัวแปรที่ใช้, สถิติที่รายงาน”
 - ถ้าเป็นกราฟ/ภาพ ใส่ “แนวโน้ม, ความสัมพันธ์, การเปรียบเทียบ”
- (3) การอ้างอิง (Attribution)
 - ถ้าเป็นข้อมูลหรือภาพที่ผู้วิจัยสร้างขึ้นเอง → ไม่ต้องอ้างที่มา
 - ถ้าคัดลอก/ดัดแปลงจากผู้อื่น → ต้องระบุแหล่งที่มา
(เช่น “ที่มา: สุภัทรชัย สีสะไบ, 2568”)
 - ใช้มาตรฐาน APA Style เช่น (สุภัทรชัย สีสะไบ, 2568)

สรุปได้ว่า การนำเสนอผลลัพธ์ที่ดีควรใช้ข้อความอธิบายความหมายของผลเสริมด้วยตารางเพื่อแสดงตัวเลขอย่างเป็นระบบ ใช้กราฟ/แผนภาพเพื่อให้เห็นแนวโน้มและการเปรียบเทียบชัดเจน และปฏิบัติตามมาตรฐาน (เช่น APA Style) เพื่อความเป็นสากล

7. การเขียนรายงานผลการวิเคราะห์

หัวข้อนี้สรุป “วิธีเล่าผลลัพธ์” จากโปรแกรมสถิติ (SPSS/JASP/R/PSPP) ให้เป็น รายงานเชิงวิชาการแบบ APA ที่อ่านแล้วเข้าใจทันทีว่า “เกิดอะไรขึ้นกับสมมติฐาน/คำถามวิจัย” ไม่ใช่เพียงคัดลอกตัวเลขจาก Output

1) โครงสร้างการรายงานผล (ที่ควรมีทุกครั้ง)

1. ระบุวัตถุประสงค์/สมมติฐาน ที่การวิเคราะห์นี้ตอบ
2. บอกสถิติที่ใช้ (เช่น independent-samples t-test, one-way ANOVA, Pearson's r, multiple regression, χ^2 test)
3. ตัวเลขสำคัญ (ครบถ้วนและเรียงลำดับ): ค่าเฉลี่ย (M), ส่วนเบี่ยงเบนมาตรฐาน (SD), ขนาดตัวอย่าง (N), ค่าตัวสถิติ (t/F/ χ^2 /r/B), อิสระ (df), ค่า p, Effect size, ช่วงความเชื่อมั่น (CI)
4. การตรวจสอบสมมติฐานเบื้องต้น/ทางเลือกเมื่อสมมติฐานไม่ผ่าน (เช่น รายงาน Welch แทน t ปกติ, ใช้ Spearman แทน Pearson)

5. การแปลความหมาย เชื่อมโยงกลับสมมติฐาน/ทฤษฎี
6. อ้างอิงตาราง/แผนภาพ ที่แนบ (“ดูตารางที่ ... / แผนภาพที่ ...”)
7. ความโปร่งใส (เช่น วิธีจัดการ Missing, เกณฑ์ outlier, การแก้ไขหลายการทดสอบ)

เคล็ดลับ เขียน “จากกว้างสู่แคบ” เริ่มจากผลโดยรวม → ตัวเลขหลัก → ผลย่อย/ post hoc → ความหมายเชิงทฤษฎี/ปฏิบัติ

2) มาตรฐานรูปแบบ (APA สั้น ๆ)

- รายงานค่า p เป็นทศนิยม 3 ตำแหน่ง $p = .028$, $p < .001$ (ไม่เขียน .000)
- ค่าเฉลี่ยและ SD ปกติ 2 ตำแหน่ง: $M = 41.87$, $SD = 9.36$
- ค่าตัวสถิติพร้อม df: $t(103.6) = -2.22$, $F(2, 57) = 10.85$, $\chi^2(2, N = 100) = 12.56$
- ใส่ Effect size ทุกครั้งที่ทำได้: Cohen’s d , η^2 /partial η^2 , r , R^2 , OR
- ใส่ช่วงความเชื่อมั่น (95% CI) เมื่อเป็นไปได้
- ระบุวิธี post hoc/การแก้ไขหลายการทดสอบ (Tukey, Bonferroni, Holm) เมื่อใช้ ANOVA หลายกลุ่ม

3) แม่แบบรายงานผล (พร้อมตัวอย่าง)

A) t-test

1) Independent-samples t-test

รายงาน: ค่าเฉลี่ยแต่ละกลุ่ม + SD, t , df , p , Cohen’s d , 95% CI ของผลต่าง, Levene’s test/ Welch (ถ้าจำเป็น)

ตัวอย่าง (เมื่อสมมติฐานความแปรปรวนเท่ากัน “ไม่ผ่าน” → ใช้ Welch)

นักเรียนหญิง ($\bar{X} = 67.00$, $SD = 2.77$, $n = 30$) มีคะแนนสูงกว่านักเรียนชาย ($\bar{X} = 57.63$, $SD = 3.90$, $n = 30$) อย่างมีนัยสำคัญทางสถิติ, Welch’s $t(\approx 58) = -11.31$, $p < .001$, Cohen’s $d = 2.89$, 95% CI ของผลต่าง [-11.04, -7.70]

ตารางที่ 8.33 Group Statistics

Gender	N	Mean (M)	SD
Male	30	57.63	3.90
Female	30	67.00	2.77

ตารางที่ 8.34 Independent Samples Test (Welch's t-test)

Test	t	df	p	Cohen's d	95% CI (Diff)
Welch's t-test	-11.31	58	<.001	2.89	[-11.04, -7.70]

2) Paired-samples t-test

คะแนนหลังเรียน ($\bar{X}$ = 62.10, SD = 3.95) สูงกว่าก่อนเรียน ($\bar{X}$ = 54.20, SD = 4.12) อย่างมีนัยสำคัญ, $t(29) = -12.35$, $p < .001$, Cohen's d (paired) = 2.25, 95% CI [-9.20, -6.60].

ตารางที่ 8.35 Paired Samples Statistics

Measure	N	Mean (M)	SD
Pre-test	30	54.20	4.12
Post-test	30	62.10	3.95

ตารางที่ 8.36 Paired Samples Test

Test	t	df	p	Cohen's d (paired)	95% CI (Diff)
Pre-test vs Post-test	-12.35	29	<.001	2.25	[-9.20, -6.60]

3) One-sample t-test

คะแนนเฉลี่ยสูงกว่าเกณฑ์ 50 อย่างมีนัยสำคัญ $t(29) = 10.75$, $p < .001$, 95% CI ของผลต่าง [6.16, 9.10]

ตารางที่ 8.37 One-sample Statistics

Variable	N	Mean (M)	SD	Test value
Test Score	30	56.16	4.25	50

ตารางที่ 8.38 One-sample Test

Test	t	df	p	95% CI (Diff)
Score vs Test value=50	10.75	29	<.001	[6

B) ANOVA

1) One-way ANOVA

รายงาน: ค่าเฉลี่ยแต่ละกลุ่ม, F , df ระหว่าง/ภายใน, p , η^2 (หรือ partial η^2), ผล post hoc ระบุวิธี

ตัวอย่าง ค่าเฉลี่ยคะแนนแตกต่างกันตามระดับการศึกษา, $F(2, 57) = 10.85$, $p < .001$, $\eta^2 = .28$. Post hoc (Tukey) พบว่า “สูง > ต่ำ” ($p < .001$) และ “กลาง > ต่ำ” ($p = .012$) ขณะที่ “สูง > กลาง” แตกต่างเล็กน้อย ($p = .045$)

ตารางที่ 8.39 Descriptive Statistics (ตามระดับการศึกษา)

Education Level	N	Mean (M)	SD
Low	20	55.20	4.10
Medium	20	59.35	3.95
High	20	62.80	4.25

ตารางที่ 8.40 ANOVA Summary

Source	SS	df	MS	F	p	η^2
Between Groups	250.2	2	125.1	10.85	<.001	.28
Within Groups	657.1	57	11.5			
Total	907.3	59				

ตารางที่ 8.41 Post hoc (Tukey HSD)

Comparison	Mean Diff	p
High – Low	7.60	<.001
Medium – Low	4.15	.012
High – Medium	3.45	.045

2) Two-way ANOVA (มีปฏิสัมพันธ์)

รายงาน: Main effects ทั้งสองตัว, Interaction, ค่า F , df , p , partial η^2 , กราฟ Interaction

ตัวอย่าง พบผลหลักของเพศ ($F(1,56) = 12.40, p = .001$, partial $\eta^2 = .18$) และวิธีสอน ($F(1,56) = 30.22, p < .001$, partial $\eta^2 = .35$) แต่ไม่พบปฏิสัมพันธ์ เพศ×วิธีสอน ($F(1,56) = 2.15, p = .148$)

ตารางที่ 8.42 Tests of Between-Subjects Effects

Source	SS	df	MS	F	p	partial η^2
Gender	120.4	1	120.4	12.40	.001	.18
Teaching Method	293.5	1	293.5	30.22	<.001	.35
Gender × Method	20.9	1	20.9	2.15	.148	.04
Error	543.1	56	9.7			

C) สหสัมพันธ์ (Correlation)

1) Pearson's r (ข้อมูลช่วง/อัตราส่วนและเป็นปกติ)

ชั่วโมงอ่านหนังสือสัมพันธ์เชิงบวกกับคะแนนหลังเรียน, $r(98) = .53, p < .001$, 95% CI [.37, .65]

ตารางที่ 8.43 Pearson Correlation Matrix

Variables	Score_Post	Hours
Score_Post	1.00	.53**
Hours	.53**	1.00

2) Spearman's ρ (ข้อมูลลำดับ/ไม่เป็นปกติ/มี outlier)

ความพึงพอใจสัมพันธ์เชิงลบกับอันดับผลสัมฤทธิ์, $\rho(98) = -.42, p < .001$

ตารางที่ 8.44 Spearman's Correlation Matrix

Variables	Satisfaction	Achievement Rank
Satisfaction	1.00	-.42**
Achievement	-.42**	1.00

หมายเหตุ: บอกเหตุผลที่ใช้ Spearman (เช่น การแจกแจงไม่เป็นปกติ)

D) การถดถอย (Regression)

1) Simple linear regression

รายงาน: F, p, R^2 , ค่าสัมประสิทธิ์ (B/β), SE, t, p , 95% CI, สมการโมเดล, ตรวจสอบสมมติฐาน (linearity, homoscedasticity, normal residuals)

ตัวอย่าง โมเดลมีนัยสำคัญ, $F(1, 98) = 102.5, p < .001, R^2 = .51$. ชั่วโมงอ่านทำนายนคะแนน ($B = 2.50, SE = 0.25, \beta = .92, t = 10.12, p < .001, 95\% \text{ CI } [2.01, 2.99]$); สมการ: $\text{Score} = 40.12 + 2.50(\text{Hours})$

ตารางที่ 8.45 Model Summary

Model	R	R ²	Adjusted R ²	F	p
1	.71	.51	.50	102.5	<.001

ตารางที่ 8.46 Coefficients (Simple Regression)

Predictor	B	SE	β	t	p	95% CI [LL, UL]
Constant	40.12	1.80	–	22.29	<.001	[36.55, 43.69]
Hours	2.50	0.25	.92	10.12	<.001	[2.01, 2.99]

2) Multiple regression

รายงาน: F, p, R^2 และ Adjusted R^2 , ค่าสัมประสิทธิ์ ทุกตัว, ตรวจสอบ multicollinearity (VIF/Tolerance), สมมติฐาน, ถ้ามีตัวแปรควบคุมให้ระบุชุด

ตัวอย่าง โมเดลหลายตัวทำนายมีนัยสำคัญ, $F(2, 97) = 56.9, p < .001, R^2 = .42$. ชั่วโมงอ่าน ($B = 1.80, \beta = .62, p = .003, \text{VIF} = 1.3$) และคะแนนก่อนเรียน ($B = 0.60, \beta = .55, p = .005, \text{VIF} = 1.3$) มีอิทธิพลอย่างมีนัยสำคัญต่อคะแนนหลังเรียน

ตารางที่ 8.47 Model Summary

Model	R	R ²	Adjusted R ²	F	p
1	.65	.42	.40	56.9	<.001

ตารางที่ 8.48 Coefficients (Multiple Regression)

Predictor	B	SE	β	t	p	VIF
Constant	35.10	2.10	–	16.71	<.001	–
Hours	1.80	0.32	.62	5.62	.003	1.3
Score_Pre	0.60	0.14	.55	4.20	.005	1.3

E) ไคสแควร์ (Chi-square test of independence)

รายงาน: Pearson χ^2 , df, p , ขนาดอิทธิพล (Cramér's V), สรุปแนวโน้มจากตารางครอสแท็บ

ตัวอย่าง เพศสัมพันธ์กับการเลือกวิชาเรียน, $\chi^2(2, N = 100) = 12.56$, $p = .002$, Cramér's $V = .35$. เพศหญิงเลือกวิชา A สูงกว่า ขณะที่เพศชายเลือกวิชา B สูงกว่า

ตารางที่ 8.49 Crosstab เพศ × การเลือกวิชาเรียน

Gender	Subject A	Subject B	Subject C	Total
Male	10	15	5	30
Female	20	7	3	30
Total	30	22	8	60

ตารางที่ 8.50 Chi-square Tests

Test	χ^2	df	p	Cramér's V
Pearson Chi-Sq.	12.56	2	.002	.35

4) การรายงานการตรวจสอบสมมติฐาน/ความทนทานของผล

- ความเป็นปกติ รายงานวิธีตรวจ (Shapiro-Wilk/กราฟ Q-Q/ดู skew-kurtosis) และผลสรุป

- ความเท่ากันของแปรปรวน รายงาน Levene's test; หากไม่ผ่านให้ใช้ Welch's t / Welch ANOVA

- อิทธิพล outlier ระบุเกณฑ์ (เช่น ± 3 SD, Cook's distance) และผลการวิเคราะห์ซ้ำเมื่อเอา outlier ออก

- หลายการทดสอบ ระบุวิธีแก้ไข (Bonferroni/Holm)

- Missing data ระบุวิธี (listwise/pairwise/MI) และอัตรา missing

ตัวอย่างถ้อยคำสั้น ๆ: “ตรวจสอบสมมติฐานไม่พบการละเมิดที่รุนแรง (Shapiro-Wilk $p > .05$; Levene's $p = .27$); ผลลัพธ์มีความสอดคล้องแม้ควบคุม outlier (Cook's $D < 1$).”

5) การอ้างอิงตาราง/แผนภาพและ Caption

- อ้างในเนื้อหา “ดังแสดงในตารางที่ 8.1” / “ดูแผนภาพที่ 8.1”

- Caption กระชับ ชัดว่าภาพ/ตารางแสดงอะไร และถ้ามีสัญลักษณ์ให้ทำ legend สั้น ๆ

- ตัวอย่าง Caption

- ตารางที่ 8.1 สรุปค่าเฉลี่ย ($\bar{X}$) และ SD ของคะแนนตามกลุ่ม พร้อม ผล t-test และ Cohen's d

- แผนภาพที่ 8.1 Interaction Plot เพศ $\times$ วิธีสอน ต่อคะแนนหลังเรียน (เส้นเกือบขนาน $\rightarrow$ ไม่มีปฏิสัมพันธ์)

6) ตัวอย่าง “ย่อหน้าแบบพร้อมรอบคิด” (Template เต็มช่องว่าง)

(ข้อสถิติ) สำหรับ (ตัวแปร/กลุ่ม) พบว่า ...

(ตัวเลขหลัก) $\bar{X}$ /SD/ N; t/F/ χ^2 / r/ R^2 , df, p, Effect size, 95% CI ...

(Assumptions/ทางเลือก) Levene/Shapiro/วิธีที่ใช้เมื่อไม่ผ่าน ...

(ความหมายเชิงวิชาการ) ผลสนับสนุน/ไม่สนับสนุนสมมติฐาน ... เชื่อมโยงทฤษฎี/งานก่อนหน้า ...

(ตาราง/แผนภาพ) ดูตารางที่ ... / แผนภาพที่ ...

ตัวอย่าง (เต็มแล้ว)

One-way ANOVA สำหรับคะแนนตามระดับการศึกษา พบความแตกต่างอย่างมีนัยสำคัญ, $F(2, 57) = 10.85, p < .001, \eta^2 = .28$. Post hoc (Tukey) ชี้ว่ากลุ่มสูง > ต่ำ ($p < .001$) และกลาง > ต่ำ ($p = .012$) ผลนี้สอดคล้องกับแนวคิดที่ว่าทักษะพื้นฐานทางวิชาการสูงสัมพันธ์กับผลสัมฤทธิ์ที่ดีกว่า (ดูตารางที่ 8.2/ แผนภาพที่ 8.2)

7) ข้อผิดพลาดที่พบบ่อย (หลีกเลี่ยง)

- เขียน $p = .000$ (ควร $p < .001$)
- สลับ SD กับ SE (ถ้าเป็น error bar ต้องระบุว่า SD หรือ SE)
- รายงานตัวเลขมากเกินไป/ทศนิยมเยอะเกินไป $\rightarrow$ อ่านยาก
- ไม่ใส่ Effect size/CI
- ไม่กล่าวถึง สมมติฐานเบื้องต้น และทางเลือกเมื่อไม่ผ่าน
- ตาราง/แผนภาพไม่มีเลขกำกับ + Caption
- รายงานผลมากมายแต่ ไม่เชื่อมโยงสมมติฐาน/บริบทวิจัย

8) เช็กลิสต์ก่อนส่งรายงาน

- ☐ ระบุสมมติฐาน/วัตถุประสงค์ที่ผลนี้ตอบ
- ☐ สถิติที่ใช้ตรงกับชนิดข้อมูลและคำถามวิจัย
- ☐ ตัวเลขครบ: M, SD, N, t/F/ χ^2 /r/ R^2 , df, p, Effect size, 95% CI
- ☐ รายงาน/จัดการ Assumptions อย่างโปร่งใส
- ☐ ใส่ตาราง/แผนภาพที่ช่วยอ่าน พร้อมเลขกำกับ + Caption

- ภาษากระชับ เเชิงความหมาย ไม่คัดลอก Output ดิบ
- รูปแบบ APA (ทศนิยม/สัญกรณ์/การอ้างอิง)
- อธิบายความหมายเชิงทฤษฎี/เชิงปฏิบัติ และข้อจำกัด

8. ข้อควรระวังในการใช้โปรแกรมทางสถิติ

การใช้โปรแกรมทางสถิติ เช่น SPSS, PSPP, JASP หรือ R ช่วยให้นักวิเคราะห์ข้อมูลสะดวกรวดเร็วและมีประสิทธิภาพ แต่ผู้วิจัยต้องตระหนักว่า โปรแกรมเป็นเพียง “เครื่องมือ” หากใช้ผิดขั้นตอน อาจทำให้ผลวิเคราะห์บิดเบือนหรือสรุปผลผิดพลาด ดังนั้นจึงควรพิจารณาประเด็นสำคัญต่อไปนี้

1) ความถูกต้องของข้อมูล (Data Accuracy)

- ต้องตรวจสอบข้อมูลก่อนนำเข้าสู่การวิเคราะห์ เช่น ค่าผิดปกติ (Outliers), Missing values, ค่าที่กรอกผิดพลาด (Data entry error)
- หากข้อมูลไม่ถูกต้อง การันสถิติใด ๆ ก็ไม่สามารถแก้ไขปัญหาที่ต้นเหตุได้
- ตัวอย่างเช่น หากมีนักเรียนอายุ “250 ปี” อยู่ใน Dataset โปรแกรมจะนำไปคำนวณจริง และทำให้ค่าเฉลี่ยบิดเบือน

2) การเลือกสถิติให้เหมาะสม (Appropriateness of Statistical Test)

- ต้องเลือกใช้สถิติตามประเภทของตัวแปรและสมมติฐาน
 - ข้อมูลเชิงจัดประเภท (Nominal/Ordinal) → ใช้ Chi-square, Mann-Whitney
 - ข้อมูลเชิงปริมาณ (Interval/Ratio) และปกติ → ใช้ t-test, ANOVA, Regression
- หากเลือกสถิติไม่ตรงกับข้อมูล อาจทำให้ผลไม่ถูกต้อง เช่น ใช้ Pearson's r กับข้อมูล Ordinal

3) การตรวจสอบสมมติฐาน (Assumption Checking)

- โปรแกรมส่วนใหญ่ไม่เตือนผู้ใชหากสมมติฐานของสถิติถูกละเมิด
- ผู้วิจัยต้องตรวจสอบเอง เช่น
 - ความเป็นปกติ (Normality) → Shapiro-Wilk, Histogram, Q-Q Plot
 - ความเท่ากันของความแปรปรวน (Homogeneity of variance) → Levene's test
 - ความเป็นอิสระของตัวอย่าง (Independence) → พิจารณาจากการออกแบบการวิจัย

- หากสมมติฐานไม่ผ่าน ควรเลือกใช้สถิติทางเลือก เช่น Welch's t-test, Kruskal-Wallis

4) การตีความผล (Interpretation of Results)

- ไม่ควรตีความจาก p-value เพียงอย่างเดียว แต่ควรดู Effect size และ Confidence interval ร่วมด้วย

- ระวังการตีความเกินกว่าขอบเขต เช่น ผล Correlation ไม่ได้หมายถึง “สาเหตุ” (Correlation $\neq$ Causation)

- รายงานผลควรอยู่ในเชิงวิชาการ ไม่ใช่เพียงบอกว่า “มี/ไม่มีนัยสำคัญ” แต่ต้องเชื่อมโยงกับสมมติฐาน ทฤษฎี และงานวิจัยก่อนหน้า

5) ความรู้ของผู้ใช้ (User Awareness)

- โปรแกรมมีเมนูสำเร็จรูป แต่ไม่ได้หมายความว่าผลลัพธ์จะต้องเสมอหากผู้ใช้ไม่เข้าใจหลักการ

- ตัวอย่างเช่น การรัน Regression ได้ค่า R^2 แต่ผู้วิจัยต้องเข้าใจว่าหมายถึง “สัดส่วนความแปรปรวนที่อธิบายได้” ไม่ใช่ “ค่าความถูกต้อง 100%”

6) การรายงานผล (Reporting)

- ผลลัพธ์ต้องนำเสนอให้ถูกต้องตามมาตรฐาน เช่น APA Style หรือรูปแบบที่สาขาวิชากำหนด

- หลีกเลี่ยงการ “ตัดต่อ” Output จนผู้อ่านเข้าใจผิด

- ควรระบุข้อมูลที่จำเป็น เช่น ค่าเฉลี่ย, SD, ค่า test statistic, df, p-value, effect size และ CI

7) การจัดการไฟล์และความสามารถของโปรแกรม

- SPSS เป็นเชิงพาณิชย์ ต้องมีลิขสิทธิ์ อาจมีข้อจำกัดด้านการเข้าถึง

- PSPP แม้ฟรี แต่รองรับการวิเคราะห์ขั้นสูงน้อยกว่า SPSS

- JASP ใช้ง่ายและฟรี แต่มีข้อจำกัดในบางฟังก์ชัน

- R ทรงพลังและยืดหยุ่น แต่ต้องมีทักษะเขียนโค้ด

สรุปได้ว่า ข้อควรระวังสำคัญคือ อย่าให้โปรแกรมเป็นผู้ตัดสินใจแทนผู้วิจัย โปรแกรมทำหน้าที่เพียงเครื่องคำนวณและแสดงผล แต่การตัดสินใจ เลือกสถิติ ตรวจสอบสมมติฐาน และตีความเชิงวิชาการเป็นความรับผิดชอบของผู้วิจัยเอง การใช้โปรแกรมอย่างมีความรู้จะทำให้งานวิเคราะห์ข้อมูลมีความถูกต้อง น่าเชื่อถือ และมีคุณค่าทางวิชาการ

สรุปท้ายบท

บทนี้ได้อธิบายถึงความสำคัญของการใช้โปรแกรมทางสถิติในการวิเคราะห์ข้อมูล โดยแนะนำโปรแกรมหลักที่นิยมใช้ ได้แก่ SPSS, JASP, R และ PSPP พร้อมทั้งอธิบายจุดเด่น ข้อจำกัด และตัวอย่างการใช้งานจริง จากนั้นได้นำเสนอหลักการป้อนข้อมูลและการจัดโครงสร้าง Dataset การตรวจสอบคุณภาพข้อมูล การวิเคราะห์เชิงพรรณนาและอนุมาน รวมทั้งตัวอย่าง Output การแปลผล และการเขียนรายงานผลตามมาตรฐาน APA

นอกจากนี้ ยังได้เน้นข้อควรระวังสำคัญ เช่น ความถูกต้องของข้อมูล การเลือกสถิติให้เหมาะสม การตรวจสอบสมมติฐาน และการตีความผลลัพธ์เชิงความหมาย เพื่อป้องกันการใช้โปรแกรมอย่างผิดพลาด ซึ่งอาจนำไปสู่การสรุปผลวิจัยที่คลาดเคลื่อน

โดยสรุป การใช้โปรแกรมทางสถิติไม่ใช่เพียงเครื่องมือคำนวณ แต่เป็นกระบวนการที่ต้องอาศัยความเข้าใจเชิงทฤษฎี ความรอบคอบในการเลือกใช้สถิติ และความสามารถในการตีความผลลัพธ์ เพื่อสร้างงานวิจัยที่มีคุณภาพ น่าเชื่อถือ และเป็นประโยชน์ต่อวงวิชาการและการประยุกต์ใช้จริง

บทที่ 9

สถิติและการวิจัยทางรัฐประศาสนศาสตร์

(STATISTICS AND RESEARCH IN PUBLIC ADMINISTRATION)

สถิติและการวิจัยเป็นรากฐานสำคัญของการศึกษาทางรัฐประศาสนศาสตร์ เนื่องจากทั้งสองเป็นเครื่องมือที่ช่วยสร้างความเข้าใจต่อปรากฏการณ์ทางการเมือง การบริหาร และสังคมอย่างมีระบบ ในอดีต การบริหารภาครัฐมักพึ่งพาประสบการณ์หรือดุลยพินิจของผู้บริหารเป็นหลัก แต่เมื่อบริบททางเศรษฐกิจ สังคม และการเมืองมีความซับซ้อนมากขึ้น การตัดสินใจเชิงนโยบายจึงต้องอาศัยข้อมูลเชิงประจักษ์และการวิเคราะห์เชิงวิชาการเป็นกลไกสนับสนุน เพื่อให้ได้มาซึ่งข้อสรุปที่มีความน่าเชื่อถือและตรวจสอบได้

ในมิติของสถิติ เครื่องมือนี้ทำหน้าที่รวบรวม จัดระเบียบ และวิเคราะห์ข้อมูลจำนวนมากมหาศาลให้อยู่ในรูปแบบที่เข้าใจง่ายและสามารถนำไปใช้ในการตัดสินใจได้อย่างมีเหตุผล ทั้งในระดับนโยบายและการจัดการภาครัฐ ขณะที่การวิจัยทำหน้าที่สร้างองค์ความรู้ใหม่ ตรวจสอบทฤษฎี และหาคำตอบต่อคำถามทางวิชาการและเชิงปฏิบัติ ซึ่งล้วนมีความสำคัญต่อการพัฒนาคุณภาพของการบริหารงานและการให้บริการสาธารณะ นอกจากนี้ การวิจัยยังทำให้เกิดการสะท้อนเสียงของประชาชนเข้าสู่กระบวนการกำหนดนโยบาย ช่วยให้การบริหารงานของรัฐบาลมีความโปร่งใสและตอบสนองต่อสังคมได้ดียิ่งขึ้น

ในยุคดิจิทัลที่การเปลี่ยนแปลงเกิดขึ้นอย่างรวดเร็ว การประยุกต์ใช้สถิติและการวิจัยในรัฐประศาสนศาสตร์ยิ่งทวีความสำคัญมากขึ้น การจัดการข้อมูลขนาดใหญ่ (Big Data) การใช้เครื่องมือวิเคราะห์เชิงขั้นสูง และการวิจัยแบบมีส่วนร่วมกำลังกลายเป็นแนวทางใหม่ที่ช่วยยกระดับคุณภาพของการตัดสินใจเชิงนโยบายและการจัดการภาครัฐ ดังนั้น บทนี้จึงมุ่งอธิบายความหมายและความสำคัญของสถิติและการวิจัย ประเภทของสถิติและวิธีการวิจัย กระบวนการวิจัยอย่างเป็นระบบ การประยุกต์ใช้สถิติในการประเมินนโยบาย กรณีศึกษาที่เกี่ยวข้อง ตลอดจนข้อจำกัด จริยธรรม เครื่องมือสมัยใหม่ และแนวโน้มอนาคต เพื่อให้ผู้ศึกษาเข้าใจบทบาทของสถิติและการวิจัยอย่างรอบด้าน และสามารถนำไปประยุกต์ใช้ในการพัฒนางานรัฐประศาสนศาสตร์ให้สอดคล้องกับมาตรฐานสากลได้อย่างมีประสิทธิภาพและยั่งยืน

1. ความหมายและความสำคัญของสถิติและการวิจัยทางรัฐประศาสนศาสตร์

สถิติและการวิจัยถือเป็นเครื่องมือสำคัญของศาสตร์รัฐประศาสนศาสตร์ เนื่องจากทั้งสองช่วยสนับสนุนการแสวงหาความรู้และการตัดสินใจเชิงนโยบายอย่างมีระบบและมีหลักฐานเชิงประจักษ์รองรับ ในบริบทของการบริหารภาครัฐที่มีความซับซ้อน การใช้สถิติทำให้ข้อมูลจำนวนมากที่กระจัดกระจายถูกจัดระเบียบและแปลความหมายได้ อย่างเป็นระบบ ขณะเดียวกัน การวิจัยทำให้เกิดองค์ความรู้ใหม่ที่สามารถนำไปใช้พัฒนา กระบวนการบริหารงานและตอบสนองต่อความต้องการของประชาชนอย่างมีประสิทธิภาพ ดังนั้น การทำความเข้าใจความหมายและความสำคัญของสถิติและการวิจัยจึงเป็นพื้นฐานที่สำคัญต่อการศึกษารัฐประศาสนศาสตร์และการประยุกต์ใช้ในเชิงนโยบาย

1) ความหมายของสถิติ

สถิติ (Statistics) ในทางรัฐประศาสนศาสตร์สามารถอธิบายได้ในหลายมิติ ดังนี้

(1) **สถิติในฐานะข้อมูล (Statistical Data)** หมายถึง ข้อมูลเชิงปริมาณที่สะท้อนข้อเท็จจริงของปรากฏการณ์ เช่น จำนวนประชากร อัตราการเกิด หรือจำนวนผู้รับบริการจากหน่วยงานของรัฐ (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุข, 2555)

(2) **สถิติในฐานะระเบียบวิธี (Statistical Method)** หมายถึง ศาสตร์ที่ว่าด้วยกระบวนการเก็บรวบรวม การนำเสนอ การวิเคราะห์ และการตีความข้อมูล เพื่อให้ได้ข้อสรุปที่เป็นระบบและเชื่อถือได้ (Kothari, 2004)

(3) **สถิติในฐานะค่าสถิติ (Statistic)** หมายถึง ค่าตัวเลขที่คำนวณจากข้อมูลตัวอย่าง เช่น ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และค่าสหสัมพันธ์ ซึ่งใช้ในการอธิบายหรือทดสอบสมมติฐานเชิงวิชาการ

จากมิติทั้งสาม สถิติในงานรัฐประศาสนศาสตร์จึงมิได้เป็นเพียงเครื่องมือทางคณิตศาสตร์เท่านั้น แต่ยังทำหน้าที่เป็นกลไกที่สำคัญในการสร้างหลักฐานเชิงประจักษ์สำหรับการกำหนดนโยบายและการบริหารงานภาครัฐ

2) ความหมายของการวิจัย

การวิจัย (Research) หมายถึง กระบวนการศึกษาค้นคว้าอย่างมีระบบตามหลักวิธีวิทยาศาสตร์ เพื่อแสวงหาความรู้หรือคำตอบที่ถูกต้องและเชื่อถือได้ (Creswell & Creswell, 2018; Kothari, 2004)

ในมิติของรัฐประศาสนศาสตร์ การวิจัยมุ่งเน้นการสร้างองค์ความรู้ใหม่หรือการพัฒนาความรู้เดิมที่เกี่ยวข้องกับโครงสร้าง กระบวนการ และผลลัพธ์ของการบริหารงานภาครัฐ ทั้งนี้เพื่อให้การดำเนินงานของรัฐเป็นไปอย่างมีประสิทธิภาพและสอดคล้องกับความต้องการของประชาชน (ประยูร อิมิวัตร์, 2566)

3) ความสำคัญของสถิติและการวิจัยทางรัฐประศาสนศาสตร์

สถิติและการวิจัยมีความสำคัญต่อรัฐประศาสนศาสตร์ในหลายประการ ดังนี้

(1) การเสริมสร้างการตัดสินใจเชิงนโยบาย ข้อมูลเชิงสถิติและผลการวิจัยช่วยให้ผู้บริหารภาครัฐสามารถตัดสินใจบนพื้นฐานของหลักฐาน ลดอคติส่วนบุคคล และสร้างความน่าเชื่อถือให้แก่การตัดสินใจ (OECD, 2020)

(2) การเพิ่มประสิทธิภาพและประสิทธิผลของการบริหารงานภาครัฐ การใช้สถิติช่วยให้การวิเคราะห์และการจัดสรรทรัพยากรของรัฐเป็นไปอย่างมีเหตุผล ทำให้การบริการสาธารณะตรงตามความต้องการของประชาชน (Gertler et al., 2016)

(3) การประเมินผลและการตรวจสอบความโปร่งใส การวิจัยทางรัฐประศาสนศาสตร์มีบทบาทสำคัญในการประเมินผลโครงการและนโยบาย เพื่อวัดผลสัมฤทธิ์ ความคุ้มค่า และความยั่งยืนของการดำเนินงาน (OECD, 2020)

(4) การสร้างองค์ความรู้ทางวิชาการ สถิติและการวิจัยเป็นรากฐานในการพัฒนาทฤษฎี แนวคิด และแบบจำลองเชิงวิชาการใหม่ ๆ เพื่อรองรับการเปลี่ยนแปลงของสังคมและการเมือง (Creswell & Creswell, 2018)

(5) การสนับสนุนการมีส่วนร่วมของประชาชน การสำรวจความคิดเห็นและการวิจัยเชิงสังคมช่วยสะท้อนมุมมองและความต้องการของประชาชนเข้าสู่กระบวนการกำหนดนโยบาย ทำให้การบริหารงานของรัฐมีลักษณะตอบสนองต่อประชาชนมากขึ้น (ประยูร อิมิวัตร์, 2566)

สรุปได้ว่า สถิติและการวิจัยเป็นกลไกสำคัญของรัฐประศาสนศาสตร์ เนื่องจากสถิติช่วยแปลงข้อมูลดิบให้เป็นสารสนเทศที่ใช้ในการวิเคราะห์และตัดสินใจ ขณะที่การวิจัยเป็นกระบวนการแสวงหาความรู้ที่มีระบบและเชื่อถือได้ การบูรณาการสถิติและการวิจัยเข้าด้วยกันจึงทำให้การบริหารงานภาครัฐเป็นไปอย่างมีประสิทธิภาพ โปร่งใส และสามารถตอบสนองต่อความต้องการของประชาชนได้อย่างแท้จริง

2. ประเภทของสถิติและวิธีการวิจัยทางรัฐประศาสนศาสตร์

การวิจัยทางรัฐประศาสนศาสตร์จำเป็นต้องอาศัยทั้งสถิติและวิธีการวิจัยที่เหมาะสม เพื่อให้ได้ข้อมูลและข้อสรุปที่มีความน่าเชื่อถือและสามารถนำไปใช้ในการบริหารงานภาครัฐได้อย่างมีประสิทธิภาพ การเลือกใช้สถิติและวิธีการวิจัยขึ้นอยู่กับลักษณะของข้อมูล วัตถุประสงค์การศึกษา และกรอบแนวคิดที่ใช้ การทำความเข้าใจประเภทของสถิติและวิธีการวิจัยจึงเป็นสิ่งสำคัญที่จะช่วยยกระดับคุณภาพงานวิจัยให้มีมาตรฐานและสามารถตอบโจทย์เชิงนโยบายได้อย่างแท้จริง

1) ประเภทของสถิติ

การใช้สถิติในงานวิจัยทางรัฐประศาสนศาสตร์สามารถจำแนกได้เป็นสองกลุ่มใหญ่ ได้แก่ สถิติเชิงพรรณนาและสถิติเชิงอนุมาน (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุโข, 2555)

(1) สถิติเชิงพรรณนา (Descriptive Statistics)

สถิติประเภทนี้ใช้เพื่อสรุปและแสดงลักษณะของข้อมูลให้ง่ายต่อการทำความเข้าใจ เช่น การใช้ค่าเฉลี่ย ร้อยละ ส่วนเบี่ยงเบนมาตรฐาน หรือการสร้างตารางแจกแจงความถี่และกราฟ (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุโข, 2555) โดยทั่วไปใช้เพื่อนำเสนอภาพรวมของข้อมูลและใช้เป็นขั้นตอนพื้นฐานก่อนการวิเคราะห์เชิงลึก

(2) สถิติเชิงอนุมาน (Inferential Statistics)

สถิติประเภทนี้ใช้เพื่ออนุมานหรือสรุปจากกลุ่มตัวอย่างไปยังประชากร โดยมุ่งเน้นการทดสอบสมมติฐาน การวิเคราะห์ความแตกต่างระหว่างกลุ่ม และการตรวจสอบความสัมพันธ์ของตัวแปร เช่น การทดสอบ t-test, การวิเคราะห์ความแปรปรวน (ANOVA), การทดสอบไคสแควร์ และการวิเคราะห์การถดถอย (Kothari, 2004; Creswell & Creswell, 2018)

ทั้งนี้ การเลือกใช้สถิติขึ้นอยู่กับลักษณะของตัวแปร ระดับการวัด และวัตถุประสงค์ของการวิจัย หากข้อมูลอยู่ในระดับนามบัญญัติหรือลำดับนิยมมักใช้สถิติอนุพารามетริก ขณะที่ข้อมูลระดับอันดับหรืออัตราส่วนมักใช้สถิติพารามетริก (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุโข, 2555)

2) ประเภทของวิธีการวิจัยทางรัฐประศาสนศาสตร์

งานวิจัยในรัฐประศาสนศาสตร์มีความหลากหลายตามบริบทและปัญหาที่ศึกษา สามารถจำแนกได้ดังนี้

(1) การวิจัยเชิงปริมาณ (Quantitative Research)

เป็นการวิจัยที่อาศัยข้อมูลเชิงตัวเลข ใช้สถิติเป็นเครื่องมือหลักในการวิเคราะห์เพื่อทดสอบสมมติฐานและอธิบายความสัมพันธ์ของตัวแปร (Kothari, 2004) ตัวอย่างเช่น การสำรวจความพึงพอใจของประชาชนต่อการบริการสาธารณะขององค์กรปกครองส่วนท้องถิ่น

(2) การวิจัยเชิงคุณภาพ (Qualitative Research)

เป็นการวิจัยที่มุ่งทำความเข้าใจปรากฏการณ์ในเชิงลึก โดยเน้นการตีความความหมายและประสบการณ์ของผู้ให้ข้อมูล ใช้วิธีการเก็บข้อมูล เช่น การสัมภาษณ์เชิงลึก การสนทนากลุ่ม หรือการสังเกตการณ์ (Creswell & Creswell, 2018) ตัวอย่างเช่น การศึกษากระบวนการตัดสินใจเชิงนโยบายของผู้บริหารระดับสูง

(3) การวิจัยแบบผสมผสาน (Mixed Methods Research)

เป็นการผสมผสานจุดแข็งของทั้งการวิจัยเชิงปริมาณและเชิงคุณภาพ โดยใช้ข้อมูลเชิงตัวเลขเพื่ออธิบายภาพรวม และใช้ข้อมูลเชิงคุณภาพเพื่อขยายความและทำความเข้าใจในเชิงลึกมากขึ้น (Creswell & Creswell, 2018) ตัวอย่างเช่น การประเมินผลโครงการพัฒนาชุมชนที่ใช้ทั้งการสำรวจเชิงปริมาณและการสัมภาษณ์ผู้มีส่วนได้ส่วนเสีย

3) ความสัมพันธ์ระหว่างสถิติและวิธีการวิจัย

สถิติและวิธีการวิจัยมีความสัมพันธ์ที่เกื้อหนุนกัน โดยวิธีการวิจัยเป็นกรอบในการออกแบบการศึกษา ขณะที่สถิติทำหน้าที่เป็นเครื่องมือในการวิเคราะห์และตีความข้อมูล การเลือกใช้สถิติที่ถูกต้องต้องสอดคล้องกับวัตถุประสงค์ของวิจัย ประเภทของข้อมูล และระดับการวัดของตัวแปร ซึ่งหากเลือกใช้ไม่เหมาะสมอาจทำให้ผลการวิจัยขาดความน่าเชื่อถือ (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุโข, 2555)

สรุปได้ว่า การเลือกใช้สถิติและวิธีการวิจัยทางรัฐประศาสนศาสตร์เป็นปัจจัยสำคัญที่ส่งผลต่อคุณภาพของงานวิจัย สถิติพรรณนาช่วยให้เห็นภาพรวมของข้อมูล ขณะที่สถิติอนุมานทำให้สามารถสรุปผลไปยังประชากรได้ ส่วนวิธีการวิจัยทั้งเชิงปริมาณเชิงคุณภาพ และแบบผสมผสานช่วยให้การศึกษาในรัฐประศาสนศาสตร์มีความสมบูรณ์ทั้งในด้านความลึกและความกว้างของข้อมูล

3. กระบวนการวิจัยทางรัฐประศาสนศาสตร์

กระบวนการวิจัยทางรัฐประศาสนศาสตร์มีความสำคัญอย่างยิ่งต่อการสร้างองค์ความรู้ที่สามารถนำไปใช้พัฒนาการบริหารงานภาครัฐและนโยบายสาธารณะ เนื่องจากเป็นการแสวงหาความรู้ที่มีระบบ ตรวจสอบได้ และมีหลักฐานเชิงประจักษ์รองรับ กระบวนการวิจัยจึงต้องดำเนินไปตามขั้นตอนที่ชัดเจนและเป็นมาตรฐาน เพื่อให้ได้ผลการวิจัยที่มีคุณภาพและสามารถอ้างอิงได้ในเชิงวิชาการและเชิงปฏิบัติ

1) การกำหนดปัญหาและวัตถุประสงค์การวิจัย

กระบวนการวิจัยเริ่มต้นด้วยการกำหนดปัญหาวิจัยอย่างชัดเจน เนื่องจากปัญหาวิจัยเป็นตัวกำหนดกรอบและทิศทางของงานวิจัยทั้งหมด (Creswell & Creswell, 2018) หากการกำหนดปัญหาไม่ชัดเจนจะทำให้การวิจัยขาดเป้าหมายและไม่สามารถตอบสนองต่อวัตถุประสงค์ที่ตั้งไว้ได้ ในบริบทของรัฐประศาสนศาสตร์ ปัญหาวิจัยมักเกี่ยวข้องกับประสิทธิภาพของการให้บริการสาธารณะ คุณภาพการบริหารงานบุคคล หรือผลกระทบของนโยบายสาธารณะต่อประชาชน

2) การทบทวนวรรณกรรมและกรอบแนวคิด

การทบทวนวรรณกรรม (Literature Review) เป็นขั้นตอนที่ช่วยให้ผู้วิจัยเข้าใจสถานภาพขององค์ความรู้ในประเด็นที่ศึกษา รวมทั้งเป็นพื้นฐานในการสร้างกรอบแนวคิดและสมมติฐาน (Kothari, 2004) ในงานรัฐประศาสนศาสตร์ การทบทวนวรรณกรรมไม่เพียงแต่พิจารณาทฤษฎีด้านการบริหารหรือการเมืองเท่านั้น แต่ยังรวมถึงงานวิจัยเชิงประจักษ์ทั้งในและต่างประเทศ เพื่อให้การวิจัยมีความครอบคลุมและเชื่อมโยงกับแนวโน้มทางวิชาการและการปฏิบัติจริง

3) การออกแบบการวิจัย

การออกแบบการวิจัย (Research Design) หมายถึง การกำหนดวิธีการและขั้นตอนอย่างเป็นระบบเพื่อให้ได้มาซึ่งข้อมูลที่ตรงตามวัตถุประสงค์ (Creswell & Creswell, 2018) การออกแบบที่เหมาะสมช่วยให้ผลการวิจัยมีความน่าเชื่อถือและสามารถตรวจสอบได้ รูปแบบที่ใช้ในรัฐประศาสนศาสตร์มีทั้งเชิงปริมาณ เชิงคุณภาพ และการวิจัยแบบผสมผสาน โดยรูปแบบที่เลือกจะขึ้นอยู่กับคำถามวิจัย ลักษณะของข้อมูล และทรัพยากรที่มีอยู่

4) การกำหนดประชากรและกลุ่มตัวอย่าง

การวิจัยภาครัฐมักเกี่ยวข้องกับประชากรจำนวนมาก ดังนั้นการเลือกกลุ่มตัวอย่างที่เป็นตัวแทนจึงมีความสำคัญอย่างยิ่ง เครื่องมือที่ใช้กันทั่วไป ได้แก่ สูตรของ Yamane (1967) และตารางกำหนดขนาดกลุ่มตัวอย่างของ Krejcie และ Morgan (1970) ซึ่งได้รับการอ้างอิงอย่างแพร่หลายในงานวิจัยทางสังคมศาสตร์และรัฐประศาสนศาสตร์ เพื่อให้มั่นใจว่าผลลัพธ์สามารถอนุมานไปยังประชากรได้อย่างเหมาะสม

5) การสร้างเครื่องมือและการเก็บรวบรวมข้อมูล

เครื่องมือวิจัยที่ใช้กันบ่อยในรัฐประศาสนศาสตร์ ได้แก่ แบบสอบถาม แบบสัมภาษณ์ และการสังเกต โดยเครื่องมือเหล่านี้ต้องผ่านการตรวจสอบความตรง (Validity) และความเชื่อมั่น (Reliability) เพื่อให้มั่นใจว่าข้อมูลที่ได้มีคุณภาพ (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุขโข, 2555) นอกจากนี้ ผู้วิจัยต้องคำนึงถึงหลักจริยธรรม เช่น การคุ้มครองสิทธิและความเป็นส่วนตัวของผู้ให้ข้อมูล

6) การวิเคราะห์ข้อมูล

ข้อมูลที่ได้จะต้องถูกวิเคราะห์ตามลักษณะของวิธีการวิจัย หากเป็นข้อมูลเชิงปริมาณจะใช้สถิติเชิงพรรณนาและสถิติเชิงอนุมาน เช่น t-test, ANOVA, Regression Analysis ส่วนข้อมูลเชิงคุณภาพจะใช้การวิเคราะห์เนื้อหา (Content Analysis) หรือการวิเคราะห์เชิงธีม (Thematic Analysis) (Creswell & Creswell, 2018)

7) การอภิปรายผลและการสรุป

ขั้นตอนสุดท้ายคือการอภิปรายผลโดยเชื่อมโยงกับกรอบแนวคิด ทฤษฎี และงานวิจัยที่เกี่ยวข้อง รวมทั้งเสนอข้อค้นพบเชิงวิชาการและข้อเสนอเชิงนโยบาย การสรุปควรชี้ให้เห็นถึงข้อจำกัดของงานวิจัยและแนวทางสำหรับการศึกษาต่อไป เพื่อเสริมสร้างองค์ความรู้ในสาขารัฐประศาสนศาสตร์อย่างต่อเนื่อง

สรุปได้ว่า กระบวนการวิจัยทางรัฐประศาสนศาสตร์ประกอบด้วยขั้นตอนที่เป็นระบบ ตั้งแต่การกำหนดปัญหา การทบทวนวรรณกรรม การออกแบบวิธีการ การกำหนดกลุ่มตัวอย่าง การสร้างเครื่องมือและเก็บข้อมูล การวิเคราะห์ ไปจนถึงการอภิปรายผลและสรุป ทั้งนี้ ความเข้มแข็งของกระบวนการวิจัยจะเป็นปัจจัยสำคัญที่ทำให้งานวิจัยสามารถสร้างองค์ความรู้ที่มีคุณภาพและนำไปใช้ในการกำหนดนโยบายและปรับปรุงการบริหารงานภาครัฐได้อย่างมีประสิทธิภาพ

4. การประยุกต์ใช้สถิติในการวิจัยและการประเมินนโยบายสาธารณะ

การประเมินนโยบายและโครงการภาครัฐถือเป็นกลไกสำคัญในการพัฒนาระบบบริหารจัดการภาครัฐให้มีความโปร่งใส มีประสิทธิภาพ และตอบสนองต่อความต้องการของประชาชน การประยุกต์ใช้สถิติในขั้นตอนดังกล่าวจึงมีบทบาทสำคัญ เนื่องจากสถิติช่วยให้การวิเคราะห์ข้อมูลมีความน่าเชื่อถือ สามารถตรวจสอบได้ และสะท้อนผลกระทบของนโยบายในเชิงประจักษ์ การนำสถิติไปใช้ในกระบวนการประเมินยังช่วยให้ผู้กำหนดนโยบายสามารถตัดสินใจได้อย่างมีเหตุผลและเป็นระบบ

1) การประเมินผลนโยบายและโครงการ

การประเมินนโยบายสาธารณะจำเป็นต้องใช้สถิติเป็นเครื่องมือหลักในการตรวจสอบผลลัพธ์และประสิทธิผลของการดำเนินงาน ทั้งในเชิงปริมาณและเชิงคุณภาพ โดยวิธีการที่นิยม ได้แก่ การทดสอบสมมติฐานเชิงสถิติ การวิเคราะห์แนวโน้ม (Trend Analysis) และการสร้างแบบจำลองทางเศรษฐมิติ (Econometric Models) เพื่อหาความสัมพันธ์ระหว่างนโยบายกับผลลัพธ์ที่เกิดขึ้น (Gertler et al., 2016) งานวิจัยเชิงประเมินผลที่มีการออกแบบการวิจัยเชิงทดลองและกึ่งทดลอง เช่น Randomized Controlled Trials (RCTs) หรือ Difference-in-Differences (DiD) มักได้รับการยอมรับว่าให้ข้อสรุปที่มีความน่าเชื่อถือสูง (Shadish, Cook, & Campbell, 2002)

2) การวิเคราะห์ต้นทุนและผลประโยชน์

การวิเคราะห์ต้นทุนและผลประโยชน์เป็นวิธีการสำคัญในการประเมินความคุ้มค่าของโครงการ โดยการเปรียบเทียบผลประโยชน์ (Benefits) ที่สังคมได้รับกับต้นทุน (Costs) ที่เกิดขึ้น การวิเคราะห์นี้อาศัยสถิติในการวัดและคำนวณ เช่น มูลค่าปัจจุบันสุทธิ

(Net Present Value: NPV), อัตราผลตอบแทนภายใน (Internal Rate of Return: IRR) และอัตราส่วนผลประโยชน์ต่อค่าใช้จ่าย (Benefit–Cost Ratio: BCR) (Boardman, Greenberg, Vining, & Weimer, 2018) การใช้สถิติในการวิเคราะห์ต้นทุนและผลประโยชน์ช่วยให้ผู้กำหนดนโยบายสามารถตัดสินใจได้อย่างมีประสิทธิภาพและมีเหตุผลและโปร่งใส โดยเฉพาะอย่างยิ่งในสภาพแวดล้อมที่ทรัพยากรมีจำกัด

3) การวิเคราะห์ประสิทธิภาพและประสิทธิผล

สถิติช่วยให้สามารถวัดระดับประสิทธิภาพ (Efficiency) และประสิทธิผล (Effectiveness) ของการดำเนินนโยบายหรือโครงการได้อย่างเป็นระบบ เช่น การใช้การวิเคราะห์เชิงถดถอย (Regression Analysis) เพื่อตรวจสอบปัจจัยที่มีผลต่อคุณภาพการบริการสาธารณะ หรือการใช้ Data Envelopment Analysis (DEA) เพื่อเปรียบเทียบประสิทธิภาพเชิงเทคนิคของหน่วยงานรัฐ (Kothari, 2004)

4) การวิเคราะห์แนวโน้มและการพยากรณ์

การใช้สถิติในการวิเคราะห์แนวโน้มและการพยากรณ์มีบทบาทสำคัญต่อการกำหนดนโยบายในอนาคต ตัวอย่างเช่น การวิเคราะห์อนุกรมเวลา (Time Series Analysis) เพื่อตรวจสอบการเปลี่ยนแปลงของอัตราการว่างงานหรือการคาดการณ์จำนวนผู้สูงอายุในอนาคต การวิเคราะห์ลักษณะนี้ช่วยให้รัฐบาลสามารถเตรียมการเชิงนโยบายล่วงหน้าได้อย่างเหมาะสม (Hyndman & Athanasopoulos, 2021)

5) การใช้สถิติในกรอบเกณฑ์สากลเพื่อการประเมิน

องค์การเพื่อความร่วมมือทางเศรษฐกิจและการพัฒนา (OECD) กำหนดเกณฑ์การประเมินที่ได้รับการยอมรับในระดับสากล ได้แก่ ความสอดคล้อง (Relevance) ความเชื่อมโยง (Coherence) ผลสัมฤทธิ์ (Effectiveness) ประสิทธิภาพ (Efficiency) ผลกระทบ (Impact) และความยั่งยืน (Sustainability) (OECD, 2020) การใช้สถิติในกรอบนี้มีบทบาทสำคัญในการสร้างมาตรฐานการประเมินที่โปร่งใสและเปรียบเทียบได้

สรุปได้ว่า การประยุกต์ใช้สถิติในการวิจัยและการประเมินนโยบายสาธารณะเป็นกลไกสำคัญที่ทำให้งานวิจัยมีความน่าเชื่อถือและสามารถนำไปใช้ประโยชน์ได้จริง สถิติทำหน้าที่ทั้งในฐานะเครื่องมือวิเคราะห์เชิงปริมาณและฐานข้อมูลเชิงประจักษ์ที่ช่วยให้การตัดสินใจเชิงนโยบายมีความโปร่งใส มีเหตุผล และตอบสนองต่อความต้องการของประชาชน การบูรณาการสถิติเข้ากับกรอบการประเมินสากลยังช่วยยกระดับคุณภาพของการบริหารงานภาครัฐให้สอดคล้องกับมาตรฐานสากล

5. กรณีศึกษาและตัวอย่างการใช้สถิติและการวิจัยทางรัฐประศาสนศาสตร์

การใช้สถิติและการวิจัยทางรัฐประศาสนศาสตร์มีได้จำกัดอยู่เพียงเชิงทฤษฎี หากแต่ปรากฏให้เห็นอย่างเป็นรูปธรรมในกรณีศึกษาต่าง ๆ ที่เกี่ยวข้องกับการบริหารงานภาครัฐและการกำหนดนโยบายสาธารณะ ตัวอย่างเหล่านี้สะท้อนให้เห็นว่าสถิติและการวิจัยสามารถนำมาใช้เพื่อวัดผล ประเมิน วิเคราะห์ และพยากรณ์ เพื่อสนับสนุนการตัดสินใจเชิงนโยบายที่มีประสิทธิภาพ โปร่งใส และตอบสนองต่อประชาชนได้จริง

1) การประเมินความพึงพอใจของประชาชนต่อการบริการสาธารณะ

งานวิจัยด้านรัฐประศาสนศาสตร์จำนวนมากมุ่งวัดระดับความพึงพอใจของประชาชนต่อการบริการของหน่วยงานภาครัฐ โดยอาศัยสถิติเชิงพรรณนา เช่น ค่าเฉลี่ยและร้อยละ เพื่อแสดงระดับความพึงพอใจโดยรวม และใช้สถิติเชิงอนุมาน เช่น t-test หรือ ANOVA เพื่อตรวจสอบความแตกต่างของความพึงพอใจระหว่างกลุ่มประชากร เช่น เพศ อายุ หรือ สถานภาพทางเศรษฐกิจ (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุโข, 2555)

2) การวิเคราะห์ปัจจัยที่มีผลต่อการมีส่วนร่วมทางการเมืองของประชาชน

การวิจัยในหัวข้อนี้มักใช้การวิเคราะห์การถดถอยเชิงพหุคูณ (Multiple Regression Analysis) เพื่อตรวจสอบว่าปัจจัยด้านการศึกษา รายได้ หรือการเข้าถึงสื่อมวลชนมีอิทธิพลต่อการมีส่วนร่วมทางการเมืองหรือไม่ (Kothari, 2004) วิธีการนี้ช่วยให้นักวิจัยสามารถแยกแยะอิทธิพลของแต่ละปัจจัยและนำไปใช้ในการออกแบบนโยบายที่ส่งเสริมการมีส่วนร่วมทางการเมืองอย่างมีประสิทธิภาพ

3) การประเมินผลโครงการพัฒนาท้องถิ่น

กรณีศึกษาการประเมินโครงการพัฒนาท้องถิ่นมักใช้วิธีการกึ่งทดลอง (Quasi-Experimental Designs) เช่น Difference-in-Differences (DiD) เพื่อเปรียบเทียบความแตกต่างของตัวชี้วัดก่อนและหลังดำเนินโครงการ ระหว่างพื้นที่ที่ได้รับและไม่ได้รับการสนับสนุน (Gertler, Martinez, Premand, Rawlings, & Vermeersch, 2016) วิธีการดังกล่าวช่วยให้ผลการประเมินมีความน่าเชื่อถือและสะท้อนผลกระทบที่แท้จริงของนโยบายต่อชุมชน

4) การวิเคราะห์ต้นทุน-ผลประโยชน์ในโครงการลงทุนภาครัฐ

การตัดสินใจลงทุนในโครงการสาธารณะ เช่น ระบบขนส่งมวลชนหรือโครงการพลังงานสะอาด มักอาศัยการวิเคราะห์ต้นทุน-ผลประโยชน์ (Cost-Benefit Analysis: CBA) โดยใช้ตัวชี้วัดทางสถิติ เช่น มูลค่าปัจจุบันสุทธิ (NPV) อัตราผลตอบแทนภายใน (IRR) และอัตราส่วนผลประโยชน์-ต้นทุน (BCR) เพื่อประเมินความคุ้มค่าของโครงการ (Boardman, Greenberg, Vining, & Weimer, 2018)

5) การพยากรณ์ความต้องการบริการสาธารณะ

หน่วยงานภาครัฐมักใช้การวิเคราะห์อนุกรมเวลา (Time Series Analysis) เพื่อคาดการณ์ความต้องการบริการในอนาคต เช่น การคาดการณ์จำนวนผู้ป่วยในโรงพยาบาลของรัฐ หรือจำนวนผู้โดยสารในระบบขนส่งสาธารณะ (Hyndman & Athanasopoulos, 2021) การใช้แบบจำลอง ARIMA หรือ Exponential Smoothing ช่วยให้การจัดสรรทรัพยากรมีประสิทธิภาพมากขึ้น และลดปัญหาความไม่สมดุลของอุปสงค์และอุปทานของบริการ

สรุปได้ว่า กรณีศึกษาข้างต้นแสดงให้เห็นว่าการใช้สถิติและการวิจัยมีความสำคัญต่อการพัฒนางานรัฐประศาสนศาสตร์ในหลายมิติ ตั้งแต่การสำรวจความคิดเห็นของประชาชน การวิเคราะห์ปัจจัยทางสังคมและการเมือง การประเมินผลโครงการพัฒนาท้องถิ่น การวิเคราะห์ความคุ้มค่าของโครงการ ไปจนถึงการพยากรณ์ความต้องการบริการสาธารณะ ทั้งหมดนี้สะท้อนว่าสถิติและการวิจัยเป็นกลไกสำคัญในการสร้างหลักฐานเชิงประจักษ์เพื่อสนับสนุนการตัดสินใจเชิงนโยบายและการปรับปรุงการบริหารงานภาครัฐอย่างมีประสิทธิภาพ

6. ข้อจำกัดของสถิติและการวิจัยทางรัฐประศาสนศาสตร์

แม้ว่าสถิติและการวิจัยจะเป็นเครื่องมือสำคัญที่ช่วยสร้างองค์ความรู้และสนับสนุนการตัดสินใจเชิงนโยบาย แต่ในการปฏิบัติจริงยังมีข้อจำกัดหลายประการที่นักวิจัยและผู้กำหนดนโยบายต้องตระหนัก ข้อจำกัดเหล่านี้เกิดขึ้นได้จาก ลักษณะของข้อมูล วิธีการเก็บข้อมูล เครื่องมือทางสถิติ รวมถึงบริบททางสังคมและการเมือง ที่ส่งผลต่อคุณภาพของงานวิจัยและการตีความผลลัพธ์

1) ข้อจำกัดด้านข้อมูล

ข้อมูลที่ใช้ในการวิจัยภาครัฐมักมีปัญหา เช่น

- **ความไม่สมบูรณ์ของข้อมูล (Incomplete Data)** ข้อมูลที่เก็บจากหน่วยงานรัฐอาจไม่ครอบคลุมทุกมิติ ทำให้การวิเคราะห์ไม่สะท้อนภาพจริงทั้งหมด
- **ความล่าช้าและไม่ทันสมัย (Outdated Data)** ข้อมูลบางชุดจัดทำเป็นรายปีหรือหลายปีครั้ง ส่งผลให้ไม่สอดคล้องกับสถานการณ์ปัจจุบัน (OECD, 2020)
- **อคติของข้อมูล (Data Bias)** อาจเกิดจากการออกแบบแบบสอบถาม การเลือกกลุ่มตัวอย่าง หรือการรายงานข้อมูลโดยผู้ให้ข้อมูล

2) ข้อจำกัดด้านระเบียบวิธีและสถิติ

- **ข้อสมมติของสถิติ** วิธีทางสถิติบางประเภท เช่น การวิเคราะห์ถดถอย หรือ ANOVA ต้องอาศัยสมมติฐานด้านการแจกแจงหรือความเป็นอิสระของตัวแปร หากข้อมูลไม่เป็นไปตามเงื่อนไข ผลลัพธ์อาจคลาดเคลื่อน (Kothari, 2004)

- **การเลือกใช้สถิติไม่เหมาะสม** นักวิจัยบางครั้งเลือกใช้สถิติที่ซับซ้อนเกินไปหรือน้อยเกินไป ซึ่งทำให้การวิเคราะห์ไม่ตอบโจทย์วัตถุประสงค์ที่แท้จริง

- **ความยากในการวัดตัวแปรทางสังคม** ตัวแปรเชิงนามธรรม เช่น ความโปร่งใส หรือการมีส่วนร่วมของประชาชน มักวัดได้ยากและต้องอาศัยการตีความ (Creswell & Creswell, 2018)

3) ข้อจำกัดด้านบริบทการวิจัย

- **ข้อจำกัดด้านงบประมาณและเวลา** งานวิจัยในภาครัฐมักมีกรอบเวลาและงบประมาณจำกัด ส่งผลให้การเก็บข้อมูลไม่ครอบคลุมเพียงพอ (Gertler et al., 2016)

- **ข้อจำกัดด้านการเข้าถึงประชากรกลุ่มเป้าหมาย** โดยเฉพาะกลุ่มที่มีสถานะทางสังคมเปราะบาง เช่น ผู้ยากไร้หรือชนกลุ่มน้อย ทำให้ผลการวิจัยอาจไม่สะท้อนภาพรวมประชากรทั้งหมด

- **อิทธิพลทางการเมืองและนโยบาย** ในบางกรณี งานวิจัยอาจถูกแทรกแซงโดยผู้มีอำนาจทางการเมือง ทำให้ผลการวิจัยไม่สามารถเผยแพร่ได้อย่างอิสระ (OECD, 2020)

4) ข้อจำกัดด้านการตีความและการนำไปใช้

- **การตีความเกินขอบเขต** บางงานวิจัยนำผลจากกลุ่มตัวอย่างเล็กไปสรุปในระดับประเทศ ซึ่งอาจไม่ถูกต้อง

- **การใช้ผลการวิจัยไม่ตรงตามเจตนา** ผู้กำหนดนโยบายอาจเลือกใช้เฉพาะผลการวิจัยที่สนับสนุนแนวนโยบายที่ต้องการ (Selective Use of Research)

- **ความไม่ต่อเนื่องของการใช้ผลการวิจัย** แม้งานวิจัยจะมีคุณภาพ แต่หากไม่มีการนำไปใช้จริง ก็ไม่ก่อให้เกิดประโยชน์เชิงนโยบาย (Boardman, Greenberg, Vining, & Weimer, 2018)

สรุปได้ว่า ข้อจำกัดของสถิติและการวิจัยทางรัฐประศาสนศาสตร์สะท้อนให้เห็นว่า การใช้เครื่องมือทางวิชาการต้องอาศัยความระมัดระวัง ไม่ว่าจะเป็นด้านข้อมูล ระเบียบวิธี บริบท และการตีความ การตระหนักถึงข้อจำกัดเหล่านี้จะช่วยให้ นักวิจัยสามารถออกแบบงานวิจัยได้รอบคอบยิ่งขึ้น และช่วยให้ผู้กำหนดนโยบายใช้ผลการวิจัยอย่างเหมาะสมและมีความรับผิดชอบ

7. จริยธรรมในการวิจัยทางรัฐประศาสนศาสตร์

การวิจัยทางรัฐประศาสนศาสตร์มิได้มุ่งเพียงสร้างองค์ความรู้เพื่อพัฒนาการบริหารงานภาครัฐ แต่ยังเกี่ยวพันโดยตรงกับชีวิต สิทธิ และศักดิ์ศรีของประชาชนผู้มีส่วนได้ส่วนเสีย ดังนั้น นักวิจัยจำเป็นต้องยึดถือจริยธรรมในการวิจัย (Research Ethics) อย่างเคร่งครัด เพื่อป้องกันการละเมิดสิทธิ ลดอคติ และสร้างความเชื่อมั่นต่อสังคมว่าผลงานวิจัยมีความโปร่งใสและน่าเชื่อถือ

1) การเคารพสิทธิและศักดิ์ศรีของผู้ให้ข้อมูล

นักวิจัยต้องเคารพในศักดิ์ศรีและสิทธิขั้นพื้นฐานของผู้ให้ข้อมูล โดยเฉพาะในเรื่องความสมัครใจและความยินยอมอย่างแจ้งชัด (Informed Consent) ผู้ให้ข้อมูลต้องได้รับคำอธิบายเกี่ยวกับวัตถุประสงค์ วิธีการ ความเสี่ยง และสิทธิในการถอนตัวออกจาก การวิจัยได้ทุกเมื่อ (Creswell & Creswell, 2018) การบังคับหรือหลอกลวงถือเป็นการละเมิดจริยธรรมที่ร้ายแรง

2) การปกป้องความลับและความเป็นส่วนตัว

การวิจัยในประเด็นที่เกี่ยวข้องกับนโยบายสาธารณะหรือข้อมูลส่วนบุคคล เช่น รายได้ ความคิดเห็นทางการเมือง หรือข้อมูลสุขภาพ ต้องคุ้มครองความลับและไม่เปิดเผยตัวตนของผู้ให้ข้อมูล เว้นแต่ได้รับความยินยอมอย่างชัดแจ้ง (Kothari, 2004) การเข้ารหัสข้อมูล (Data Encryption) และการใช้รหัสแทนชื่อจริงเป็นแนวทางที่ช่วยรักษาความปลอดภัยของข้อมูล

3) ความซื่อสัตย์และความโปร่งใสในการวิจัย

นักวิจัยต้องดำเนินงานด้วยความซื่อสัตย์ ทั้งในการเก็บข้อมูล การวิเคราะห์ และการรายงานผล ห้ามบิดเบือนหรือสร้างข้อมูลขึ้นเอง (Fabrication and Falsification) รวมถึงต้องอ้างอิงแหล่งที่มาของแนวคิดหรือผลงานวิชาการอย่างถูกต้องตามหลักสากล (APA, 2020) ความโปร่งใсыังหมายรวมถึงการเปิดเผยข้อจำกัดของงานวิจัย ไม่เลือกนำเสนอเฉพาะผลลัพธ์ที่สนับสนุนสมมติฐานของตน

4) การหลีกเลี่ยงผลประโยชน์ทับซ้อน

งานวิจัยทางรัฐประศาสนศาสตร์มักเกี่ยวข้องกับนโยบายและทรัพยากรของรัฐ นักวิจัยจึงต้องระมัดระวังไม่ให้ผลประโยชน์ส่วนตัวหรือสถาบันที่สนับสนุนงานวิจัยมีอิทธิพลต่อผลลัพธ์และข้อเสนอเชิงนโยบาย (OECD, 2020) การเปิดเผยแหล่งทุนและความสัมพันธ์กับหน่วยงานที่เกี่ยวข้องถือเป็นมาตรฐานที่ช่วยสร้างความโปร่งใส

5) ความรับผิดชอบต่อสังคม

การวิจัยภาครัฐควรมุ่งสร้างประโยชน์ต่อส่วนรวม ไม่ก่อให้เกิดผลเสียหรือความขัดแย้งในสังคม นักวิจัยต้องใช้ผลการวิจัยเพื่อสนับสนุนการพัฒนาที่เป็นธรรมและยั่งยืน

(Gertler, Martinez, Premand, Rawlings, & Vermeersch, 2016) อีกทั้งควรสื่อสารผลการวิจัยต่อสาธารณะในลักษณะที่เข้าใจง่าย เพื่อเสริมสร้างการมีส่วนร่วมของประชาชน

สรุปได้ว่า จริยธรรมในการวิจัยทางรัฐประศาสนศาสตร์ครอบคลุมทั้งการเคารพสิทธิและศักดิ์ศรีของผู้ให้ข้อมูล การปกป้องความลับ ความซื่อสัตย์ ความโปร่งใส การหลีกเลี่ยงผลประโยชน์ทับซ้อน และความรับผิดชอบต่อสังคม การยึดมั่นในหลักจริยธรรมเหล่านี้ไม่เพียงช่วยคุ้มครองผู้มีส่วนได้ส่วนเสีย แต่ยังช่วยยกระดับคุณภาพและความน่าเชื่อถือของงานวิจัย และทำให้ผลงานมีคุณภาพต่อการพัฒนาระบบบริหารงานภาครัฐอย่างแท้จริง

8. เครื่องมือและเทคโนโลยีสมัยใหม่ในการวิจัยทางรัฐประศาสนศาสตร์

การเปลี่ยนแปลงทางเทคโนโลยีในศตวรรษที่ 21 ส่งผลโดยตรงต่อวิธีการและเครื่องมือที่ใช้ในการวิจัยทางรัฐประศาสนศาสตร์ เทคโนโลยีสารสนเทศและดิจิทัลไม่เพียงเพิ่มประสิทธิภาพในการเก็บและประมวลผลข้อมูล แต่ยังเปิดโอกาสให้นักวิจัยเข้าถึงข้อมูลขนาดใหญ่ (Big Data) วิเคราะห์ด้วยเครื่องมือที่ซับซ้อน และสื่อสารผลการวิจัยต่อสาธารณะได้อย่างรวดเร็วและกว้างขวาง การทำความเข้าใจเครื่องมือและเทคโนโลยีเหล่านี้จึงเป็นสิ่งจำเป็นสำหรับการพัฒนางานวิจัยภาครัฐให้สอดคล้องกับยุคสมัย

1) โปรแกรมสถิติและซอฟต์แวร์วิเคราะห์ข้อมูล

การวิจัยทางรัฐประศาสนศาสตร์อาศัยโปรแกรมคอมพิวเตอร์ในการจัดการและวิเคราะห์ข้อมูลจำนวนมาก เช่น

- SPSS และ Stata สำหรับการวิเคราะห์เชิงสถิติพื้นฐานและขั้นสูง
- R และ Python สำหรับการเขียนโค้ดและการสร้างแบบจำลองที่ซับซ้อน รวมถึงการวิเคราะห์ข้อมูลเชิงพยากรณ์

- NVivo และ ATLAS.ti สำหรับการวิเคราะห์ข้อมูลเชิงคุณภาพ เช่น การถอดรหัสการสัมภาษณ์หรือการวิเคราะห์เชิงธีม (Creswell & Creswell, 2018)

2) ฐานข้อมูลและคลังสารสนเทศ

เทคโนโลยีดิจิทัลทำให้นักวิจัยเข้าถึงฐานข้อมูลออนไลน์ทั้งในระดับประเทศและนานาชาติได้ง่ายขึ้น เช่น

- World Bank Data, OECD Data สำหรับข้อมูลเศรษฐกิจและการพัฒนา
- สำนักงานสถิติแห่งชาติ สำหรับข้อมูลประชากรและเศรษฐกิจไทย

- Scopus และ Web of Science สำหรับการสืบค้นงานวิจัยและวารสารวิชาการ การใช้ฐานข้อมูลเหล่านี้ช่วยให้การวิจัยมีความน่าเชื่อถือและเชื่อมโยงกับมาตรฐานสากล (OECD, 2020)

3) การเก็บข้อมูลออนไลน์และเทคโนโลยีสำรวจ

ปัจจุบันมีเครื่องมือดิจิทัลที่ช่วยให้การเก็บข้อมูลสะดวกและรวดเร็ว เช่น

- แบบสอบถามออนไลน์ผ่าน Google Forms, SurveyMonkey, Qualtrics
- การวิเคราะห์ความคิดเห็นจาก Social Media Analytics เช่น Twitter, Facebook หรือ TikTok

- การใช้ GIS (Geographic Information System) ในการวิเคราะห์ข้อมูลเชิงพื้นที่ เช่น การกระจายตัวของประชากรหรือการเข้าถึงบริการสาธารณะ (Bivand, Pebesma, & Gómez-Rubio, 2013)

4) การวิเคราะห์ข้อมูลขนาดใหญ่ (Big Data) และปัญญาประดิษฐ์ (AI)

การวิจัยในยุคดิจิทัลไม่จำกัดเฉพาะข้อมูลจากการสำรวจ แต่ยังรวมถึงข้อมูลขนาดใหญ่ที่ได้จากระบบดิจิทัล เช่น ข้อมูลการใช้บริการภาครัฐออนไลน์ หรือข้อมูลจากเซ็นเซอร์ในเมืองอัจฉริยะ (Smart Cities) เทคโนโลยีปัญญาประดิษฐ์ (AI) และการเรียนรู้ของเครื่อง (Machine Learning) ช่วยให้นักวิจัยสามารถค้นหารูปแบบที่ซับซ้อนและทำการพยากรณ์ได้อย่างแม่นยำ (Brynjolfsson & McAfee, 2017)

5) การสื่อสารและเผยแพร่ผลการวิจัย

การเผยแพร่ผลการวิจัยสามารถทำได้รวดเร็วและเข้าถึงผู้มีส่วนได้ส่วนเสียหลากหลายกลุ่มมากขึ้นผ่านเครื่องมือดิจิทัล เช่น

- Dashboard ออนไลน์ ที่แสดงผลข้อมูลแบบอินเตอร์แอคทีฟ
- Infographic และ Visualization Tools เช่น Tableau, Power BI ที่ช่วยให้ข้อมูลเชิงซับซ้อนเข้าใจง่าย

- Open Access Publishing เพื่อให้ผลงานวิจัยเผยแพร่สู่สาธารณะได้โดยไม่ถูกจำกัด

สรุปได้ว่า เครื่องมือและเทคโนโลยีสมัยใหม่มีบทบาทสำคัญในการเพิ่มคุณภาพและประสิทธิภาพของงานวิจัยทางรัฐประศาสนศาสตร์ ตั้งแต่การเข้าถึงฐานข้อมูล การเก็บข้อมูล การวิเคราะห์เชิงสถิติและ Big Data ไปจนถึงการเผยแพร่ผลการวิจัย เทคโนโลยีเหล่านี้ไม่เพียงช่วยให้นักวิจัยมีความทันสมัยและเชื่อถือได้ แต่ยังสนับสนุนการตัดสินใจเชิงนโยบายที่สอดคล้องกับความเปลี่ยนแปลงของโลกยุคดิจิทัล

9. แนวโน้มและทิศทางอนาคตของการวิจัยและการใช้สถิติทางรัฐประศาสนศาสตร์

การเปลี่ยนแปลงทางเศรษฐกิจ สังคม และเทคโนโลยี ได้ส่งผลให้ทิศทางของงานวิจัยและการใช้สถิติในรัฐประศาสนศาสตร์ก้าวเข้าสู่มิติใหม่ จากเดิมที่เน้นการใช้สถิติพื้นฐานเพื่ออธิบายปรากฏการณ์ ขยายไปสู่การวิเคราะห์ข้อมูลเชิงลึก การพยากรณ์แนวโน้ม และการสร้างแบบจำลองที่ซับซ้อนมากขึ้น อนาคตของงานวิจัยในสาขานี้จึงจำเป็นต้องผสมผสานทั้งองค์ความรู้ด้านทฤษฎีและการใช้เครื่องมือดิจิทัลขั้นสูง เพื่อสนับสนุนการกำหนดนโยบายที่สอดคล้องกับการเปลี่ยนแปลงของโลก

1) การใช้ Big Data และ Open Data

- ข้อมูลขนาดใหญ่จากหน่วยงานรัฐและแพลตฟอร์มดิจิทัล จะถูกนำมาใช้เพื่อวิเคราะห์พฤติกรรม ความต้องการ และแนวโน้มของประชาชน

- การเปิดเผยข้อมูลสาธารณะ (Open Data) จะเป็นแนวโน้มสำคัญเพื่อเพิ่มความโปร่งใส และเปิดโอกาสให้ภาคประชาชนและนักวิชาการร่วมตรวจสอบและใช้ประโยชน์ (Janssen et al., 2017)

2) ปัญญาประดิษฐ์และการเรียนรู้ของเครื่อง (AI & Machine Learning)

- AI จะช่วยให้การวิเคราะห์เชิงพยากรณ์แม่นยำยิ่งขึ้น เช่น การคาดการณ์ความต้องการบริการสาธารณะ หรือการตรวจสอบความเสี่ยงด้านนโยบาย

- การเรียนรู้ของเครื่องช่วยให้สามารถตรวจจับรูปแบบซ่อนเร้นของข้อมูลที่สถิติแบบดั้งเดิมอาจมองไม่เห็น (Brynjolfsson & McAfee, 2017)

3) การผสมผสานข้อมูลเชิงปริมาณและคุณภาพแบบ Real-Time

- แนวโน้มใหม่คือการเก็บและวิเคราะห์ข้อมูลแบบ **เรียลไทม์** ผ่านเซ็นเซอร์ แอปพลิเคชัน และระบบดิจิทัล

- งานวิจัยจะไม่แยกเชิงปริมาณและคุณภาพออกจากกัน แต่ใช้ข้อมูลเชิงตัวเลขเพื่อบอกภาพรวม และใช้ข้อมูลเชิงคุณภาพเพื่ออธิบายความหมายเชิงลึกอย่างต่อเนื่อง (Creswell & Creswell, 2018)

4) การวิจัยเชิงนโยบายที่มีส่วนร่วม

- แนวโน้มในอนาคตจะมุ่งเน้นการทำวิจัยที่ประชาชนและผู้มีส่วนได้ส่วนเสียเข้ามามีส่วนร่วม ตั้งแต่การกำหนดปัญหาวิจัยจนถึงการสื่อสารผลลัพธ์

- ส่งผลให้การกำหนดนโยบายตอบสนองต่อความต้องการของประชาชนมากขึ้น และสร้างความชอบธรรมในการตัดสินใจ (Fung, 2015)

5) การใช้เกณฑ์การประเมินสากลและการเปรียบเทียบระหว่างประเทศ

การประเมินนโยบายสาธารณะจะยึดตามเกณฑ์สากล เช่น OECD DAC Criteria: Relevance, Coherence, Effectiveness, Efficiency, Impact, และ Sustainability (OECD, 2020)

แนวโน้มงานวิจัยเปรียบเทียบ (Comparative Research) จะมีบทบาทสำคัญ เพื่อให้การบริหารจัดการของรัฐไทยเชื่อมโยงกับมาตรฐานสากล

สรุปได้ว่า แนวโน้มและทิศทางอนาคตของการวิจัยและการใช้สถิติในรัฐประศาสนศาสตร์ จะก้าวไปสู่การใช้ข้อมูลและเทคโนโลยีขั้นสูง เช่น Big Data, AI, และ Real-Time Analytics ผสานกับการทำวิจัยเชิงนโยบายที่มีส่วนร่วมและยึดเกณฑ์สากลมากขึ้น สิ่งเหล่านี้จะช่วยให้การบริหารจัดการภาครัฐมีความโปร่งใส มีประสิทธิภาพ และสามารถตอบสนองต่อพลวัตทางเศรษฐกิจและสังคมได้อย่างยั่งยืน

สรุปท้ายบท

สถิติและการวิจัยในทางรัฐประศาสนศาสตร์ถือเป็นกลไกสำคัญที่ทำให้ศาสตร์การบริหารงานภาครัฐมีความเป็นระบบและมีหลักฐานเชิงประจักษ์รองรับ เนื้อหาตลอดทั้งบทได้สะท้อนให้เห็นว่า สถิติทำหน้าที่ไม่เพียงแต่เป็นข้อมูลเชิงปริมาณหรือวิธีการทางคณิตศาสตร์เท่านั้น หากยังเป็นเครื่องมือที่ช่วยจัดระเบียบข้อมูลจำนวนมากให้นำไปวิเคราะห์และตีความได้อย่างมีเหตุผล ขณะที่การวิจัยทำหน้าที่สร้างองค์ความรู้ใหม่และตรวจสอบสมมติฐานที่เกี่ยวข้องกับการบริหารงานและนโยบายสาธารณะ เมื่อบูรณาการเข้าด้วยกัน ทั้งสองจึงกลายเป็นรากฐานของการตัดสินใจเชิงนโยบายที่มีคุณภาพ

การทำความเข้าใจประเภทของสถิติและวิธีการวิจัยทำให้เห็นถึงความหลากหลายของเครื่องมือที่สามารถนำมาใช้ได้อย่างเหมาะสมตามวัตถุประสงค์ ไม่ว่าจะเป็นสถิติเชิงพรรณนาที่ช่วยสรุปภาพรวมของข้อมูล สถิติเชิงอนุมานที่ช่วยให้สามารถขยายผลจากตัวอย่างไปสู่ประชากร หรือวิธีการวิจัยเชิงปริมาณ เชิงคุณภาพ และแบบผสมผสานที่เปิดโอกาสให้การศึกษาในรัฐประศาสนศาสตร์มีความสมบูรณ์ทั้งในด้านความกว้างและความลึก กระบวนการวิจัยที่เป็นระบบ ตั้งแต่การกำหนดปัญหา การทบทวนวรรณกรรม การออกแบบวิธีวิจัย การเก็บและวิเคราะห์ข้อมูล ไปจนถึงการอภิปรายและสรุปผล ล้วนมีความสำคัญต่อการสร้างองค์ความรู้ที่ตรวจสอบได้และนำไปใช้ได้จริง

ในทางปฏิบัติ การประยุกต์ใช้สถิติยังมีบทบาทชัดเจนต่อการประเมินนโยบายและโครงการสาธารณะ การวิเคราะห์ต้นทุนและผลประโยชน์ การตรวจสอบประสิทธิภาพและประสิทธิผล ตลอดจนการพยากรณ์แนวโน้มในอนาคต กรณีศึกษาต่าง ๆ เช่น การวิเคราะห์ความพึงพอใจของประชาชนต่อบริการสาธารณะ การมีส่วนร่วมทางการเมือง การ

ประเมินโครงการพัฒนาท้องถิ่น หรือการพยากรณ์ความต้องการบริการ ล้วนสะท้อนให้เห็นว่าสถิติและการวิจัยสามารถแปลงข้อเท็จจริงเชิงประจักษ์ไปสู่ข้อเสนอเชิงนโยบายได้อย่างเป็นรูปธรรม

อย่างไรก็ตาม งานวิจัยภาครัฐก็ยังมีข้อจำกัดที่ต้องตระหนัก ไม่ว่าจะเป็นความไม่สมบูรณ์ของข้อมูล ความซับซ้อนของวิธีการทางสถิติ ข้อจำกัดด้านเวลาและงบประมาณ หรืออิทธิพลทางการเมืองที่อาจกระทบต่อการตีความและการนำผลวิจัยไปใช้จริง ปัญหาเหล่านี้ทำให้การรักษารายธรรมทางวิชาการยิ่งมีความสำคัญ นักวิจัยต้องเคารพสถิติผู้ให้ข้อมูล ปกป้องความลับส่วนบุคคล ดำเนินงานด้วยความซื่อสัตย์ หลีกเลี่ยงผลประโยชน์ทับซ้อน และมุ่งรับผิดชอบต่อสังคม เพื่อให้ผลการวิจัยมีความน่าเชื่อถือและเป็นประโยชน์ต่อสาธารณะอย่างแท้จริง

นอกจากนี้ การเปลี่ยนแปลงทางเทคโนโลยียังทำให้สถิติและการวิจัยในรัฐประศาสนศาสตร์ก้าวเข้าสู่มิติใหม่ เครื่องมือดิจิทัล โปรแกรมวิเคราะห์ข้อมูล ฐานข้อมูลออนไลน์ Big Data ปัญญาประดิษฐ์ และเทคโนโลยี GIS ล้วนช่วยเพิ่มพลังให้กับการวิจัยทั้งในด้านความเร็ว ความถูกต้อง และความสามารถในการตีความข้อมูลเชิงซับซ้อน ขณะเดียวกัน แนวโน้มอนาคตยังชี้ให้เห็นการพัฒนาไปสู่การใช้ข้อมูลแบบเรียลไทม์ การผสมผสานวิธีวิจัยเชิงปริมาณและคุณภาพเข้าด้วยกัน การทำวิจัยเชิงนโยบายที่มีส่วนร่วมของประชาชน และการยึดเกณฑ์การประเมินสากลเพื่อสร้างมาตรฐานในการบริหารงานภาครัฐ

กล่าวโดยสรุป บทนี้แสดงให้เห็นว่าสถิติและการวิจัยเป็นเสาหลักที่เกื้อหนุนกัน โดยสถิติทำหน้าที่เปลี่ยนข้อมูลดิบให้เป็นสารสนเทศที่น่าไปใช้วิเคราะห์ ส่วนการวิจัยสร้างกระบวนการแสวงหาความรู้ที่มีระบบและน่าเชื่อถือ เมื่อบูรณาการเข้าด้วยกัน ทั้งสองจึงทำให้รัฐประศาสนศาสตร์ก้าวไปสู่การเป็นศาสตร์ที่มีความโปร่งใส มีเหตุผล และสามารถตอบสนองต่อความต้องการของประชาชนได้อย่างแท้จริง

บทที่ 10

ตัวอย่างการใช้สถิติในงานวิจัยทางรัฐประศาสนศาสตร์ (EXAMPLES OF THE USE OF STATISTICS IN PUBLIC ADMINISTRATION RESEARCH)

งานวิจัยทางรัฐประศาสนศาสตร์มีจุดมุ่งหมายสำคัญในการสร้างองค์ความรู้เพื่อพัฒนาการบริหารจัดการภาครัฐ ตลอดจนการกำหนดนโยบายสาธารณะที่ตอบสนองต่อความต้องการของประชาชนได้อย่างแท้จริง การที่จะทำให้งานวิจัยดังกล่าวมีความน่าเชื่อถือและสามารถนำไปใช้ได้จริงนั้น จำเป็นต้องอาศัยสถิติ เป็นเครื่องมือหลักในการจัดการ วิเคราะห์ และตีความข้อมูล การใช้สถิติในงานวิจัยไม่ได้จำกัดอยู่เพียงการอธิบายข้อมูลเบื้องต้นเท่านั้น แต่ยังครอบคลุมถึงการทดสอบสมมติฐาน การตรวจสอบความสัมพันธ์เชิงซับซ้อนระหว่างตัวแปร และการสร้างแบบจำลองทางทฤษฎีที่สะท้อนความเป็นจริงในเชิงประจักษ์ การแบ่งประเภทการใช้สถิติสามารถทำได้ 3 ระดับหลัก ได้แก่

1. สถิติเชิงพรรณนา ใช้เพื่อสรุปและนำเสนอข้อมูลพื้นฐานของกลุ่มตัวอย่าง เช่น เพศ อายุ ระดับการศึกษา รายได้ หรือความคิดเห็นเฉลี่ยของประชาชน วิธีการที่นิยมได้แก่ การหาความถี่ ร้อยละ ค่าเฉลี่ย และส่วนเบี่ยงเบนมาตรฐาน การใช้สถิติเชิงพรรณนาช่วยให้ผู้วิจัยมองเห็นภาพรวมของข้อมูลได้ชัดเจนและเป็นระบบ

2. สถิติเชิงอนุมาน เป็นการนำข้อมูลจากกลุ่มตัวอย่างมาสรุปและอ้างอิงไปยังประชากรทั้งหมด โดยอาศัยหลักความน่าจะเป็นและการทดสอบสมมติฐาน เช่น การทดสอบค่าเฉลี่ยระหว่างกลุ่ม การวิเคราะห์ความสัมพันธ์ และการวิเคราะห์การถดถอย สถิติเชิงอนุมานจึงเป็นเครื่องมือสำคัญที่ช่วยยืนยันหรือปฏิเสธสมมติฐานทางวิชาการ และทำให้ข้อค้นพบมีความน่าเชื่อถือในเชิงสถิติ

3. สถิติขั้นสูง ถูกนำมาใช้มากขึ้นในงานวิจัยร่วมสมัย โดยเฉพาะเมื่อต้องวิเคราะห์ตัวแปรที่มีหลายมิติและความสัมพันธ์เชิงซ้อน เช่น การวิเคราะห์องค์ประกอบเชิงยืนยัน (CONFIRMATORY FACTOR ANALYSIS: CFA) เพื่อยืนยันความตรงเชิงโครงสร้างของโมเดล หรือการวิเคราะห์สมการโครงสร้าง (STRUCTURAL EQUATION MODELING: SEM) เพื่อตรวจสอบความสัมพันธ์เชิงสาเหตุระหว่างตัวแปรแฝงและตัวแปรสังเกต การประยุกต์ใช้สถิติ

ชั้นสูงช่วยให้สามารถสร้างแบบจำลองเชิงทฤษฎีที่สอดคล้องกับข้อมูลเชิงประจักษ์ และยกระดับคุณภาพงานวิจัยให้เป็นที่ยอมรับในระดับสากล

ดังนั้น การใช้สถิติในงานวิจัยรัฐประศาสนศาสตร์ไม่เพียงเป็นการวิเคราะห์ข้อมูล แต่ยังเป็นการสร้างหลักฐานเชิงประจักษ์ ที่นำไปสู่การกำหนดนโยบายสาธารณะ และการบริหารจัดการภาครัฐอย่างมีเหตุผล โปร่งใส และตรวจสอบได้ ทั้งนี้ การเลือกใช้สถิติแต่ละประเภทควรพิจารณาให้เหมาะสมกับวัตถุประสงค์งานวิจัย ลักษณะของข้อมูล และกรอบแนวคิดทางทฤษฎีที่ใช้กับการศึกษา

1. การใช้สถิติเชิงพรรณนา (Descriptive Statistics)

สถิติเชิงพรรณนา (Descriptive Statistics) เป็นสถิติที่ใช้ในการสรุปผลข้อมูลที่ซับซ้อนให้ง่ายต่อความเข้าใจ โดยอาศัยการแจกแจงความถี่ ค่าเฉลี่ย ร้อยละ มัธยฐาน ฐานนิยม ส่วนเบี่ยงเบนมาตรฐาน และการวัดการกระจาย จุดประสงค์หลักคือเพื่ออธิบายลักษณะทั่วไปของข้อมูล โดยไม่ได้มุ่งอนุมานไปยังประชากร แต่ใช้เพื่อให้เห็นภาพรวมของกลุ่มตัวอย่างหรือตัวแปรที่ศึกษา (ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุขโข, 2555)

ในงานวิจัยทางรัฐประศาสนศาสตร์ สถิติเชิงพรรณนามักใช้เพื่อนำเสนอข้อมูลพื้นฐานของกลุ่มตัวอย่าง เช่น เพศ อายุ ระดับการศึกษา รายได้ รวมทั้งใช้วิเคราะห์ความคิดเห็นและทัศนคติที่วัดด้วยมาตราส่วนประมาณค่า (Rating Scale) โดยมักหาค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานเพื่อแสดงระดับการเห็นด้วยหรือการปฏิบัติของกลุ่มตัวอย่าง (Creswell & Creswell, 2018)

ตัวอย่างการใช้สถิติเชิงพรรณนาจากงานวิจัยทางรัฐประศาสนศาสตร์

งานวิจัยของพงศ์นคร โภชากรณ์ (2566) เรื่อง การบูรณาการหลักพุทธธรรมเพื่อการบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ ได้ใช้สถิติเชิงพรรณนาเป็นเครื่องมือหลักในการอธิบายลักษณะทั่วไปของกลุ่มตัวอย่าง และการตอบแบบสอบถามที่เกี่ยวข้องกับการบริหารจัดการเชิงนโยบาย โดยรวบรวมข้อมูลจากกลุ่มตัวอย่าง 372 คน และวิเคราะห์ด้วยโปรแกรมสำเร็จรูปทางสถิติ

1) ข้อมูลส่วนบุคคลของกลุ่มตัวอย่าง

สถิติเชิงพรรณนาถูกนำมาใช้เพื่อนำเสนอข้อมูลพื้นฐาน เช่น ความถี่ (Frequency) และร้อยละ (Percentage) ของคุณลักษณะส่วนบุคคลของผู้ตอบแบบสอบถาม ดังตัวอย่างเช่น ผลการวิเคราะห์ข้อมูลทั่วไปเกี่ยวกับปัจจัยส่วนบุคคลของผู้ตอบแบบสอบถาม จำนวน 372 คน จำแนกตาม เพศ อายุ ที่อยู่ปัจจุบัน ระดับการศึกษา

ความเกี่ยวข้องกับโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ และระยะเวลาที่ผ่านเกี่ยวข้องกับโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ แสดงด้วยความถี่และร้อยละมีรายละเอียดดังแสดงในตารางดังนี้

ตารางที่ 10.1 จำนวนและร้อยละของผู้ตอบแบบสอบถามจำแนกตามปัจจัยส่วนบุคคล
(n=372)

สถานภาพทั่วไปของผู้ตอบแบบสอบถาม	จำนวน	ร้อยละ
เพศ		
ชาย	82	22.04
หญิง	290	77.96
รวม	372	100.00
อายุ		
18 – 30 ปี	215	57.80
31 – 40 ปี	88	23.66
41 – 50 ปี	50	13.44
51 ปีขึ้นไป	19	5.10
รวม	372	100.00
ภูมิลำเนาปัจจุบัน (จังหวัดที่ทำงานอยู่ในปัจจุบัน)		
จังหวัดเชียงใหม่	47	12.63
จังหวัดนครราชสีมา	47	12.63
จังหวัดกาญจนบุรี	18	4.84
จังหวัดพระนครศรีอยุธยา	23	6.18
จังหวัดชลบุรี	34	9.14
จังหวัดนครศรีธรรมราช	37	9.95
กรุงเทพมหานคร	166	44.63
รวม	372	100.00
ระดับการศึกษาสูงสุด		
ต่ำกว่าปริญญาตรี	16	4.30
ปริญญาตรี	313	84.14
สูงกว่าปริญญาตรี	43	11.56
รวม	372	100.00

ตารางที่ 10.1 จำนวนและร้อยละของผู้ตอบแบบสอบถามจำแนกตามปัจจัยส่วนบุคคล (ต่อ)
(n=372)

สถานภาพทั่วไปของผู้ตอบแบบสอบถาม	จำนวน	ร้อยละ
ความเกี่ยวข้องกับโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ		
เป็นผู้บริหารที่ร่วมกำหนดทิศทาง	4	1.08
เป็นบุคลากรในหน่วยงานนโยบายหรือหน่วยงานหลัก	3	0.81
เป็นบุคลากรในหน่วยงานรับลงทะเบียน	357	95.97
เป็นบุคลากรในหน่วยงานตรวจสอบคุณสมบัติ	8	2.14
รวม	372	100.00
ระยะเวลาที่เกี่ยวข้องกับโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ		
ต่ำกว่า 6 เดือน	114	30.65
6 เดือน – 1 ปี	197	52.96
1 - 2 ปี	10	2.69
2 – 3 ปี	4	1.08
มากกว่า 3 ปี	47	12.62
รวม	372	100.00

จากตารางที่ 10.1 พบว่า ข้อมูลทั่วไปของผู้ตอบแบบสอบถามเรื่องการบูรณาการหลักพุทธธรรมเพื่อพัฒนาการบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ จำแนกได้ดังนี้

เพศ พบว่า ผู้ตอบแบบสอบถามส่วนใหญ่เป็น เพศหญิง 290 คน คิดเป็นร้อยละ 77.96 และเพศชาย 82 คน คิดเป็นร้อยละ 22.04

อายุ พบว่า ส่วนใหญ่มีอายุ 18–30 ปี จำนวน 215 คน คิดเป็นร้อยละ 57.80 รองลงมาคือ อายุ 31–40 ปี จำนวน 88 คน คิดเป็นร้อยละ 23.66 อายุ 41–50 ปี จำนวน 50 คน คิดเป็นร้อยละ 13.44 อายุ 51 ปีขึ้นไป จำนวน 19 คน คิดเป็นร้อยละ 5.10 และ อายุ 61 ปีขึ้นไป จำนวน 1 คน คิดเป็นร้อยละ 0.27

ภูมิลำเนาปัจจุบัน (จังหวัดที่ทำงาน) พบว่า มีภูมิลำเนาอยู่ในกรุงเทพมหานคร จำนวน 166 คน คิดเป็นร้อยละ 44.63 รองลงมาคือ จังหวัดเชียงใหม่ 47 คน คิดเป็นร้อยละ 12.63 จังหวัดนครราชสีมา 47 คน คิดเป็นร้อยละ 12.63 จังหวัดนครศรีธรรมราช 37 คน

คิดเป็นร้อยละ 9.95 จังหวัดชลบุรี 34 คน คิดเป็นร้อยละ 9.14 จังหวัดพระนครศรีอยุธยา 23 คน คิดเป็นร้อยละ 6.18 และจังหวัดกาญจนบุรี 18 คน คิดเป็นร้อยละ 4.84

ระดับการศึกษา พบว่า ผู้ตอบส่วนใหญ่จบการศึกษาในระดับปริญญาตรี จำนวน 313 คน คิดเป็นร้อยละ 84.14 รองลงมาคือ ระดับสูงกว่าปริญญาตรี 43 คน คิดเป็นร้อยละ 11.56 และระดับต่ำกว่าปริญญาตรี 16 คน คิดเป็นร้อยละ 4.30

ความเกี่ยวข้องกับโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ พบว่า ส่วนใหญ่เป็นบุคลากรในหน่วยงานรับลงทะเบียน 357 คน คิดเป็นร้อยละ 95.97 รองลงมาคือ บุคลากรในหน่วยงานตรวจสอบคุณสมบัติ 8 คน คิดเป็นร้อยละ 2.14 ผู้บริหารที่ร่วมกำหนดทิศทาง 4 คน คิดเป็นร้อยละ 1.08 และบุคลากรในหน่วยงานนโยบายหรือหน่วยงานหลัก 3 คน คิดเป็นร้อยละ 0.81

ระยะเวลาที่เกี่ยวข้องกับโครงการ พบว่า ส่วนใหญ่มีความเกี่ยวข้องในช่วง 6 เดือน – 1 ปี จำนวน 197 คน คิดเป็นร้อยละ 52.96 รองลงมาคือ ต่ำกว่า 6 เดือน 114 คน คิดเป็นร้อยละ 30.65 มากกว่า 3 ปี 47 คน คิดเป็นร้อยละ 12.62 1–2 ปี 10 คน คิดเป็นร้อยละ 2.69 และ 2–3 ปี 4 คน คิดเป็นร้อยละ 1.08

จากผลการวิเคราะห์ข้างต้น แสดงให้เห็นว่ากลุ่มตัวอย่างมีความหลากหลายทั้งด้านเพศ อายุ ภูมิภาค ระดับการศึกษา และระดับความเกี่ยวข้องกับโครงการ โดยใช้สถิติเชิงพรรณนา ได้แก่ ความถี่ (Frequency) และร้อยละ (Percentage) เพื่อนำเสนอในรูปแบบที่เข้าใจง่ายและเป็นระบบ (พงศ์นคร โกษากรณ์, 2566)

2) ผลการวิเคราะห์ข้อมูล

การบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ ประกอบด้วย 3 ด้านหลัก ได้แก่

1. ด้านความแม่นยำในการคัดกรองผู้มีรายได้ร้อยละตามคุณสมบัติที่กำหนด
2. ด้านการปรับปรุงฐานข้อมูลผู้มีรายได้ร้อยละให้เป็นปัจจุบัน
3. ด้านประสิทธิภาพในการจัดสรรสวัสดิการให้แก่ผู้มีรายได้ร้อยละ

ผลการวิเคราะห์โดยใช้ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐาน มีรายละเอียดดังแสดงใน ตารางที่ 10.2

ตารางที่ 10.2 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และระดับการปฏิบัติตามการบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ โดยภาพรวม

(n=372)

การบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ	ระดับการปฏิบัติ		
	$\bar{X}$	S.D.	แปลผล
1. ด้านความแม่นยำในการคัดกรองผู้มีรายได้น้อยตามคุณสมบัติที่กำหนด	3.81	0.67	มาก
2. ด้านการปรับปรุงฐานข้อมูลผู้มีรายได้น้อยให้เป็นปัจจุบัน	3.94	0.65	มาก
3. ด้านประสิทธิภาพในการจัดสรรสวัสดิการให้แก่ผู้มีรายได้น้อย	3.91	0.71	มาก
ภาพรวม	3.89	0.57	มาก

จากตารางที่ 10.2 ผลการวิเคราะห์ข้อมูลเกี่ยวกับระดับการปฏิบัติตามการบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ ภาพรวมอยู่ในระดับมาก ($\bar{X} = 3.89$, S.D. = 0.57) เมื่อพิจารณาเป็นรายด้าน พบว่าอยู่ในระดับมากทุกด้าน ได้แก่ ด้านการปรับปรุงฐานข้อมูลผู้มีรายได้น้อยให้เป็นปัจจุบัน มีค่าเฉลี่ยสูงสุดที่สุด ($\bar{X} = 3.94$, S.D. = 0.65) แสดงว่าหน่วยงานที่เกี่ยวข้องมีการปรับปรุงข้อมูลให้เป็นปัจจุบันอยู่เสมอ เพื่อใช้ประกอบการพิจารณาสิทธิได้อย่างถูกต้อง ด้านประสิทธิภาพในการจัดสรรสวัสดิการให้แก่ผู้มีรายได้น้อย มีค่าเฉลี่ยรองลงมา ($\bar{X} = 3.91$, S.D. = 0.71) สะท้อนถึงความพยายามของหน่วยงานรัฐในการส่งต่อสิทธิประโยชน์แก่ผู้มีสิทธิอย่างทั่วถึง และด้านความแม่นยำในการคัดกรองผู้มีรายได้น้อยตามคุณสมบัติที่กำหนด มีค่าเฉลี่ยน้อยที่สุด ($\bar{X} = 3.81$, S.D. = 0.67) แม้อยู่ในระดับมาก แต่สะท้อนให้เห็นว่ายังมีพื้นที่ในการพัฒนา โดยเฉพาะการเพิ่มความถูกต้องของกระบวนการคัดกรองเพื่อลดข้อผิดพลาดในการระบุผู้มีสิทธิ

กล่าวโดยสรุป ผลการวิเคราะห์แสดงให้เห็นว่า การบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐอยู่ในระดับที่น่าพอใจทุกด้าน โดยเฉพาะการปรับปรุงฐานข้อมูลและการจัดสรรสวัสดิการ แต่ยังคงเสริมมาตรการให้ความแม่นยำในการคัดกรองมีความเข้มงวดและมีประสิทธิภาพมากยิ่งขึ้น

2. การใช้สถิติเชิงอนุมาน (Inferential Statistics)

สถิติเชิงอนุมาน หมายถึง วิธีการวิเคราะห์ข้อมูลที่ใช้กลุ่มตัวอย่างเพื่อนำไปอนุมานหรือสรุปผลแทนประชากร โดยอาศัยการทดสอบสมมติฐาน การตรวจสอบความแตกต่าง และการวิเคราะห์ความสัมพันธ์ของตัวแปร สถิติเชิงอนุมานช่วยให้นักวิจัยสามารถตรวจสอบว่าผลลัพธ์ที่ได้จากการศึกษาเป็นเพียงความบังเอิญ หรือมีนัยสำคัญเชิงสถิติจริง (Creswell & Creswell, 2018)

วัตถุประสงค์ของการใช้สถิติเชิงอนุมาน

1. เพื่อทดสอบสมมติฐานที่ตั้งขึ้นว่ามีความแตกต่างหรือความสัมพันธ์ระหว่างตัวแปรหรือไม่
2. เพื่อสรุปผลจากข้อมูลตัวอย่างไปยังประชากรโดยอาศัยค่าความน่าจะเป็น
3. เพื่อใช้เป็นหลักฐานในการตัดสินใจเชิงนโยบายหรือเชิงบริหาร

การประยุกต์ใช้ในงานวิจัยรัฐประศาสนศาสตร์

งานวิจัยด้านรัฐประศาสนศาสตร์มักใช้สถิติเชิงอนุมานเพื่อตรวจสอบสมมติฐาน เช่น การศึกษาความแตกต่างของความคิดเห็นตามเพศ อายุ ระดับการศึกษา หรือสถานภาพการทำงาน ตลอดจนการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร เช่น การมีส่วนร่วมทางการเมืองกับความเชื่อมั่นต่อภาครัฐ เป็นต้น

ตัวอย่างการใช้สถิติเชิงอนุมานจากงานวิจัยทางรัฐประศาสนศาสตร์

1) การทดสอบ t-test

การทดสอบ t-test ใช้เพื่อเปรียบเทียบค่าเฉลี่ยของ 2 กลุ่ม ว่ามีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติหรือไม่ ตัวอย่างเช่น การเปรียบเทียบความคิดเห็นของบุคลากรชายและหญิงต่อคุณภาพการให้บริการสาธารณะ

ตัวอย่างงานวิจัย

งานวิจัยของชญาณุช สามัญ (2561) เรื่อง การพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอน้อย จังหวัดพระนครศรีอยุธยา ใช้ Independent Sample t-test เพื่อเปรียบเทียบการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอน้อย จังหวัดพระนครศรีอยุธยา โดยจำแนกตามเพศ อายุ ระดับการศึกษา ตำแหน่ง และระยะเวลาในการปฏิบัติงาน โดยใช้การทดสอบค่า (t-test) ในกรณีตัวแปรต้นสองกลุ่ม และการทดสอบค่าเอฟ (F-test) ด้วยวิธีการวิเคราะห์ความแปรปรวนทางเดียว (One Way Anova) ในกรณีตัวแปรต้นตั้งแต่สามกลุ่มขึ้นไป เมื่อพบว่ามีความแตกต่างจะ

ทำการเปรียบเทียบความแตกต่างรายคู่ ด้วยวิธี (Least Significant Difference: LSD.) โดยกำหนดนัยสำคัญทางสถิติที่ระดับ 0.05

สมมติฐานที่ 1 บุคลากรที่มีเพศต่างกันมีระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา แตกต่างกัน

ตารางที่ 10.3 การเปรียบเทียบระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา ภาพรวม จำแนกตามเพศ

(n=128)

การพัฒนาทรัพยากร มนุษย์	เพศ				t	Sig.
	ชาย (55 คน)		หญิง (55 คน)			
	$\bar{X}$	S.D.	$\bar{X}$	S.D.		
1. ด้านการฝึกอบรม	4.30	0.177	4.34	0.156	-1.109	0.270
2. ด้านการศึกษา	4.33	0.173	4.32	0.158	0.165	0.870
3. ด้านการพัฒนา	4.32	0.148	4.35	0.204	-0.945	0.347
รวม	4.32	0.101	4.34	0.098	-1.106	0.271

จากตารางที่ 10.3 ผลการเปรียบเทียบระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาล เมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา โดยภาพรวม จำแนกตามเพศ พบว่า โดยภาพรวมไม่แตกต่างกัน ($t = -1.106$, Sig. = 0.271) ดังนั้นจึงปฏิเสธสมมติฐานการวิจัย เมื่อพิจารณาเป็นรายด้านพบว่า ทุกด้านไม่แตกต่างกัน ดังนั้นจึงปฏิเสธสมมติฐานการวิจัย

จากผลการวิเคราะห์ดังกล่าว สรุปได้ว่า การพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา ไม่แตกต่างกันระหว่างบุคลากรเพศชายและหญิง ทั้งในภาพรวมและรายด้าน แสดงให้เห็นว่าเทศบาลมีการพัฒนาทรัพยากรมนุษย์อย่างเท่าเทียม โดยไม่จำกัดเพศของบุคลากร

2) การทดสอบ ANOVA

การวิเคราะห์ความแปรปรวน (Analysis of Variance: ANOVA) ใช้เพื่อเปรียบเทียบค่าเฉลี่ยของมากกว่า 2 กลุ่ม ว่ามีความแตกต่างกันหรือไม่

ตัวอย่างงานวิจัย

งานวิจัยของชญาณุช สามัญ (2561) เรื่อง การพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา ใช้การทดสอบค่าเอฟ

(F-test) เพื่อเปรียบเทียบการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตวน้อย จังหวัดพระนครศรีอยุธยา โดยจำแนกตามเพศ อายุ ระดับการศึกษา ตำแหน่ง และระยะเวลาในการปฏิบัติงาน โดยใช้การทดสอบค่าเอฟ (F-test) ด้วยวิธีการวิเคราะห์ความแปรปรวนทางเดียว (One Way Anova) ในกรณีตัวแปรต้นตั้งแต่สามกลุ่มขึ้นไป เมื่อพบว่ามีความแตกต่างจะทำการเปรียบเทียบความแตกต่างรายคู่ ด้วยวิธี (Least Significant Difference: LSD.) โดยกำหนดนัยสำคัญทางสถิติที่ระดับ 0.05

สมมติฐานที่ 3 บุคลากรที่มีระดับการศึกษาต่างกันมีระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตวน้อย จังหวัดพระนครศรีอยุธยา แตกต่างกัน

ตารางที่ 10.4 การเปรียบเทียบระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตวน้อย จังหวัดพระนครศรีอยุธยา ภาพรวม จำแนกตามระดับการศึกษา (n=128)

การพัฒนา ทรัพยากรมนุษย์	แหล่งความ แปรปรวน	SS	df	MS	F	Sig.
1. ด้านการฝึกอบรม	ระหว่างกลุ่ม	0.021	2	0.011	0.388	0.679
	ภายในกลุ่ม	3.463	125	0.028		
	รวม	3.484	127			
2. ด้านการศึกษา	ระหว่างกลุ่ม	0.118	2	0.059	2.229	0.112
	ภายในกลุ่ม	3.306	125	0.026		
	รวม	3.424	127			
3. ด้านการพัฒนา	ระหว่างกลุ่ม	0.502	2	0.251	8.486**	0.000
	ภายในกลุ่ม	3.698	125	0.030		
	รวม	4.200	127			
รวม	ระหว่างกลุ่ม	0.035	2	0.017	1.785	0.172
	ภายในกลุ่ม	1.223	125	0.010		
	รวม	1.258	127			

**มีนัยสำคัญทางสถิติที่ระดับ 0.01

จากตารางที่ 10.4 ผลการเปรียบเทียบระดับการพัฒนาทรัพยากรมนุษย์ของเทศบาล เมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา โดยภาพรวม จำแนกตามระดับการศึกษา พบว่า โดยภาพรวมไม่แตกต่างกัน ($F = 1.785$, $Sig. = 0.127$) ดังนั้นจึงปฏิเสธสมมติฐานการวิจัย

เมื่อพิจารณาเป็นรายด้านพบว่า ด้านการฝึกอบรมและด้านการศึกษาบุคลากรมีระดับการพัฒนาไม่แตกต่างกัน ส่วนด้านการพัฒนา มีความแตกต่างกัน ($F = 8.486$, $Sig. = 0.000$) อย่างมีนัยสำคัญที่ระดับ 0.01 ดังนั้นจึงทำการเปรียบเทียบความแตกต่างค่าเฉลี่ยเป็นรายคู่โดยวิธี ผลต่างนัยสำคัญน้อยที่สุด (Least Significant Difference: LSD.) ในด้านการพัฒนา ดังตารางที่ 10.5

ตารางที่ 10.5 การเปรียบเทียบความแตกต่างรายคู่ด้วยวิธีผลต่างนัยสำคัญน้อยที่สุดของการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา ด้านการพัฒนา จำแนกตามระดับการศึกษา

(n=128)

การพัฒนา ทรัพยากรมนุษย์		ระดับการศึกษา		
		ต่ำกว่าปริญญาตรี	ปริญญาตรี	สูงกว่าปริญญาตรี
		4.26	4.36	4.47
1. ต่ำกว่าปริญญาตรี	4.26	-	.010*	-0.21*
2. ปริญญาตรี	4.36	-	-	-0.11
3. สูงกว่าปริญญาตรี	4.47	-	-	-

*มีนัยสำคัญทางสถิติที่ระดับ 0.05

จากตารางที่ 10.5 ผลการเปรียบเทียบความแตกต่างรายคู่ด้วยวิธีผลต่างนัยสำคัญน้อยที่สุดของการพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอลำตาเสา จังหวัดพระนครศรีอยุธยา ด้านการพัฒนา จำแนกตามระดับการศึกษา พบว่า บุคลากรที่มีวุฒิการศึกษาต่ำกว่าปริญญาตรีมีระดับพัฒนาน้อยกว่าบุคลากรที่มีวุฒิการศึกษาระดับปริญญาตรี และวุฒิสูงกว่าปริญญาตรีอย่างมีนัยสำคัญที่ระดับ 0.05

ผลลัพธ์สะท้อนว่า ระดับการศึกษามีอิทธิพลต่อการพัฒนาทรัพยากรมนุษย์ในด้านการพัฒนา โดยเฉพาะผู้ที่มีการศึกษาสูงกว่ามีการพัฒนามากกว่า ซึ่งสอดคล้องกับทฤษฎีการบริหารทรัพยากรมนุษย์ที่ชี้ว่าการศึกษามีบทบาทสำคัญต่อการเพิ่มศักยภาพและการพัฒนาตนเองของบุคลากร

3) การวิเคราะห์ความสัมพันธ์ (Correlation Analysis)

การวิเคราะห์ความสัมพันธ์เป็นเทคนิคทางสถิติที่ใช้ตรวจสอบว่าตัวแปรสองตัวหรือมากกว่ามีความสัมพันธ์กันหรือไม่ และมีทิศทางอย่างไร โดยทั่วไปมักใช้ค่าสัมประสิทธิ์สหสัมพันธ์เพียร์สัน (Pearson's r) สำหรับตัวแปรที่มีการวัดในระดับอันตรภาค (Interval) หรืออัตราส่วน (Ratio) และสหสัมพันธ์สเปียร์แมน (Spearman's ρ) สำหรับตัวแปรที่มีการวัดในระดับลำดับ (Ordinal)

ค่า r มีค่าอยู่ระหว่าง -1 ถึง +1

ถ้า $r > 0$ แสดงว่ามีความสัมพันธ์เชิงบวก

ถ้า $r < 0$ แสดงว่ามีความสัมพันธ์เชิงลบ

ถ้า $r = 0$ หมายถึงไม่มีความสัมพันธ์

ขนาดของความสัมพันธ์มักตีความว่า (Cohen, 1988)

0.00–0.30 = ความสัมพันธ์ต่ำ

0.31–0.70 = ความสัมพันธ์ปานกลาง

0.71–1.00 = ความสัมพันธ์สูง

ตัวอย่างงานวิจัย

งานวิจัยของนายยา อินทร์กรุงเก่า (2566) เรื่อง ปัจจัยที่ส่งผลต่อประสิทธิภาพในการปฏิบัติงานตามหลักพุทธธรรมของบุคลากรมหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย ใช้การวิเคราะห์สหสัมพันธ์เพียร์สัน (Pearson's Correlation Coefficient) เพื่อตรวจสอบความสัมพันธ์ระหว่างตัวแปรอิสระด้านปัจจัยการปฏิบัติงานและหลักอิทธิบาท 4 กับตัวแปรตามด้านประสิทธิภาพในการปฏิบัติงานของบุคลากร

ผลการวิเคราะห์ความสัมพันธ์

ผลการวิเคราะห์ข้อมูลความสัมพันธ์ระหว่างปัจจัยการปฏิบัติงานและหลักอิทธิบาท 4 กับประสิทธิภาพในการปฏิบัติงานของบุคลากรมหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย มีรายละเอียดดังแสดงในตารางที่ 10.6

ตารางที่ 10.6 ความสัมพันธ์ระหว่างปัจจัยจิตใจหรือปัจจัยที่เป็นตัวกระตุ้นในการปฏิบัติงาน ปัจจัยคำจูนหรือปัจจัยสุขอนามัย และหลักอิทธิบาท 4 กับ ประสิทธิภาพในการปฏิบัติงานของบุคลากรมหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย

		ประสิทธิภาพในการปฏิบัติงาน ของบุคลากรมหาวิทยาลัยมหา จุฬาลงกรณราชวิทยาลัย
ปัจจัยจิตใจหรือปัจจัย ที่เป็นตัวกระตุ้นใน การปฏิบัติงาน	ค่าสหสัมพันธ์	.689**
	P	0.01
	ระดับความสัมพันธ์	ปานกลาง
ปัจจัยคำจูนหรือปัจจัย สุขอนามัย	ค่าสหสัมพันธ์	.795**
	P	0.01
	ระดับความสัมพันธ์	ปานกลาง
หลักอิทธิบาท 4	ค่าสหสัมพันธ์	.618**
	P	0.01
	ระดับความสัมพันธ์	ปานกลาง

**มีนัยสำคัญทางสถิติที่ระดับ 0.01

การตีความผลการวิจัย

จากตารางข้างต้นพบว่า

1. ปัจจัยจิตใจหรือปัจจัยที่เป็นตัวกระตุ้นในการปฏิบัติงานมีความสัมพันธ์เชิงบวกกับประสิทธิภาพการปฏิบัติงาน ($r = .689$, $p < .01$) สะท้อนว่าการมีแรงจูงใจที่เหมาะสม เช่น โอกาสก้าวหน้าและการได้รับการยอมรับ ช่วยเพิ่มประสิทธิภาพในการทำงานได้ในระดับปานกลาง
2. ปัจจัยคำจูนหรือปัจจัยสุขอนามัยมีความสัมพันธ์เชิงบวกสูงสุด ($r = .795$, $p < .01$) ซึ่งว่าการจัดสภาพแวดล้อมที่ดี เช่น ค่าตอบแทนที่เหมาะสม ความมั่นคงในการทำงาน และสภาพการทำงานที่ปลอดภัย ส่งผลต่อประสิทธิภาพได้อย่างมีนัยสำคัญ
3. หลักอิทธิบาท 4 ได้แก่ ฉันทะ วิริยะ จิตตะ และวิมังสา มีความสัมพันธ์เชิงบวกกับประสิทธิภาพการปฏิบัติงาน ($r = .618$, $p < .01$) แสดงว่าการที่บุคลากรยึด

หลักการปฏิบัติที่เน้นความเพียร ความตั้งใจ และการใช้ปัญญา สามารถช่วยยกระดับประสิทธิภาพการทำงานได้ในระดับปานกลาง

โดยสรุปงานวิจัยนี้ยืนยันว่า ทั้งปัจจัยการทำงานและหลักพุทธธรรม (อิทธิบาท 4) ล้วนมีความสัมพันธ์อย่างมีนัยสำคัญกับประสิทธิภาพการปฏิบัติงานของบุคลากรและสามารถนำไปใช้เป็นแนวทางเชิงนโยบายในการพัฒนาทรัพยากรมนุษย์ขององค์กรได้อย่างเป็นรูปธรรม (นาตยา อินทร์กรุงเก่า, 2566)

4) การวิเคราะห์การถดถอย (Regression Analysis)

การวิเคราะห์การถดถอย เป็นวิธีการทางสถิติที่ใช้เพื่อศึกษาความสัมพันธ์เชิงเหตุผล (Causal Relationship) และเชิงพยากรณ์ (Prediction) ระหว่างตัวแปรอิสระ (Independent Variables) และตัวแปรตาม (Dependent Variable) โดยมีจุดมุ่งหมายเพื่อ

1. อธิบายความสัมพันธ์ของตัวแปรในเชิงทฤษฎี
2. ทำนายค่าของตัวแปรตามจากค่าของตัวแปรอิสระ
3. ตรวจสอบอิทธิพลของตัวแปรอิสระที่มีต่อตัวแปรตาม

การวิเคราะห์ถดถอยสามารถแบ่งได้เป็น 2 ประเภทหลัก ได้แก่

1. การถดถอยเชิงเส้นอย่างง่าย (Simple Linear Regression) ใช้ตัวแปรอิสระเพียง 1 ตัว เพื่อทำนายค่าตัวแปรตาม เช่น การใช้จำนวนชั่วโมงอบรมเพื่อทำนายระดับประสิทธิภาพการทำงาน

2. การถดถอยพหุคูณ (Multiple Regression) ใช้ตัวแปรอิสระมากกว่า 1 ตัว เพื่อทำนายค่าตัวแปรตาม เช่น การใช้ปัจจัยด้านการจัดการ เวลา และงบประมาณ เพื่อทำนายผลสัมฤทธิ์โครงการ

การประยุกต์ใช้ในงานวิจัยทางรัฐประศาสนศาสตร์

งานวิจัยทางรัฐประศาสนศาสตร์ใช้การถดถอยเพื่อทดสอบว่าปัจจัยใดบ้างมีอิทธิพลต่อการบริหารงานภาครัฐ เช่น การบริหารโครงการ การพัฒนาทรัพยากรมนุษย์ หรือการมีส่วนร่วมทางการเมืองของประชาชน การถดถอยจึงเป็นเครื่องมือสำคัญในการทำตัวแปรทำนายที่แท้จริง และช่วยสร้างแบบจำลองเชิงนโยบาย (Policy Models)

ตัวอย่างงานวิจัย

งานวิจัยของ พระมหาสมัคร อติภทโท (2568) เรื่อง การบูรณาการหลักพุทธธรรมเพื่อส่งเสริมความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ใช้การวิเคราะห์การถดถอยพหุคูณแบบ Stepwise เพื่อหาปัจจัยที่ส่งผล

ต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา โดยตัวแปรอิสระ 2 ตัว คือ 1. การglomเกลตาทางการเมือง ได้แก่ 1) ด้านครอบครัว 2) ด้านกลุ่มเพื่อน 3) ด้านสถาบันการศึกษา 4) ด้านกลุ่มทางสังคม 2. หลักสังเกต 4 ได้แก่ 1) ทาน โอบอ้อมอารี 2) ปิยวาจา วจีไพเราะ 3) อัจฉริยา สงเคราะห์ประชาชน 4) สมานัตตา วางตนพอดี ตัวแปรตาม ได้แก่ ความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยาทั้ง 4 ด้าน ประกอบด้วย 1) ด้านความต่อเนื่องในการเลือกพรรค 2) ด้านการตัดสินใจล่วงหน้าในการเลือกพรรค 3) ด้านความเชื่อมั่นในการเป็นตัวแทนของพรรค 4) ด้านการสนับสนุนกิจการของพรรค

ผลการวิเคราะห์การถดถอย

ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร เพื่อใช้สร้างเมทริกซ์สหสัมพันธ์ การบูรณาการหลักพุทธธรรมเพื่อส่งเสริมความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

1) สัญลักษณ์ที่ใช้ในการวิเคราะห์ข้อมูล

t	แทน	ค่าสถิติที่ใช้พิจารณาใน t-Distribution
r	แทน	ค่าสัมประสิทธิ์สหสัมพันธ์ของเพียร์สัน (Pearson's Product Moment Correlation Coefficient)
R	แทน	ค่าสัมประสิทธิ์สหสัมพันธ์พหุคูณ
R Square	แทน	ค่าสหสัมพันธ์พหุคูณยกกำลังสอง
Std. Error	แทน	ค่าความคลาดเคลื่อนมาตรฐานในพยากรณ์
B	แทน	ค่าสัมประสิทธิ์การถดถอยพหุคูณของตัวแปรพยากรณ์ในรูปคะแนนดิบ
Beta	แทน	ค่าสัมประสิทธิ์การถดถอยพหุคูณตัวแปรพยากรณ์ในรูปคะแนนมาตรฐาน
$\hat{Y}$	แทน	ค่าสัมประสิทธิ์การถดถอยพหุคูณของตัวแปรที่ถูกพยากรณ์ในรูปคะแนนมาตรฐาน
Sig.	แทน	ระดับนัยสำคัญ

ตัวแปรการglomเกลตาทางการเมือง

SumA	แทน	ด้านครอบครัว
SumB	แทน	ด้านกลุ่มเพื่อน

SumC	แทน	ด้านสถาบันการศึกษา
------	-----	--------------------

SumD	แทน	กลุ่มทางสังคม
------	-----	---------------

ตัวแปรหลักสังคหวัด 4

SumAA	แทน	ทาน โอบอ้อมอารี
-------	-----	-----------------

SumBB	แทน	ปิยวาจา วจีไพเราะ
-------	-----	-------------------

SumCC	แทน	อัทธจริยา สงเคราะห์ประชาชน
-------	-----	----------------------------

SumDD	แทน	สมานัตตตา วางตนพอดี
-------	-----	---------------------

ตัวแปรความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองใน

จังหวัดพระนครศรีอยุธยา

SumAAA	แทน	ด้านความต่อเนื่องในการเลือกพรรค
--------	-----	---------------------------------

SumBBB	แทน	ด้านการตัดสินใจล่วงหน้าในการเลือกพรรค
--------	-----	---------------------------------------

SumCCC	แทน	ด้านความเชื่อมั่นในการเป็นตัวแทนของพรรค
--------	-----	---

SumDDD	แทน	ด้านการสนับสนุนกิจการของพรรค
--------	-----	------------------------------

Y	แทน	ผลรวม
---	-----	-------

2) ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร เพื่อใช้สร้างเมทริกซ์สหสัมพันธ์การบูรณาการหลักพุทธธรรมเพื่อส่งเสริมความ
นิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

ตารางที่ 10.7 การทดสอบสัมประสิทธิ์สหสัมพันธ์ระหว่างการglomเกลตาทางการเมืองและหลักสังคหวัตถุ 4

	SumA	SumB	SumC	SumD	รวม ABCD	SumAA	SumBB	SumCC	SumDD	Sum AABBCCDD	Y
SumA	1										
SumB	.473**	1									
SumC	-.082	-.201**	1								
SumD	.138**	.096	-.042	1							
SumABCD	.597**	.581**	.314**	.660**	1						
SumAA	.101*	.311**	-.121*	.320**	.309**	1					
SumBB	.495**	.502**	-.239**	.362**	.500**	.303**	1				
SumCC	.249**	.326**	.073	.480**	.550**	.386**	.369**	1			
SumDD	.173**	.337**	.025	.174**	.326**	.217**	.305**	.605**	1		
SumAABBCCDD	.356**	.519**	-.123*	.466**	.576**	.739**	.716**	.756**	.651**	1	
Y	.182**	.439**	.240**	.398**	.607**	.147**	.367**	.306**	.268**	.367**	1

** . Correlation is significant at the 0.01 level (2-tailed).

* . Correlation is significant at the 0.05 level (2-tailed).

จากตารางที่ 10.7 พบว่า การglomเกลตาทางการเมืองโดยส่วนใหญ่มีความสัมพันธ์ในเชิงบวกกับหลักสังคหวัตถุ 4 อย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 และ 0.05

ค่าสัมประสิทธิ์สหสัมพันธ์ระหว่างการกล่อมเกลາทางการเมืองมีสัมพันธภาพอยู่ระหว่าง (-0.239 ถึง 0.635) ถือว่าไม่มีปัญหา Collinearity & Multicollinearity จึงต้องมีการทดสอบ Collinearity Statistics โดยพิจารณาจากค่า (Tolerance) และค่าปัจจัยการขยายตัวของความแปรปรวน (Variance Inflation Factor : VIF) ดังตารางที่ 10.8

ตารางที่ 10.8 ค่า Tolerance และค่า Variance Inflation Factor (VIF) ของการกล่อมเกลາทางการเมืองและหลักสังคหวัตถุ 4

ตัวแปร	Collinearity Statistics	
	Tolerance	VIF
ตัวแปรการกล่อมเกลาทางการเมือง		
ด้านครอบครัว	0.666	1.502
ด้านกลุ่มเพื่อน	0.591	1.691
ด้านสถาบันการศึกษา	0.883	1.132
กลุ่มทางสังคม	0.664	1.506
ตัวแปรหลักสังคหวัตถุ 4		
ทาน โอบอ้อมอารี	0.664	1.506
ปิยวาจา วจีไพเราะ	0.760	1.316
อิตถจริยา สงเคราะห์ประชาชน	0.544	1.837
สมานัตตตา วางตนพอดี	0.441	2.270

จากตารางที่ 10.8 ดัชนีบอกภาวะร่วมเส้นทางตรงพหุค่าความคงทนของการยอมรับ (Tolerance) ได้ค่าระหว่าง 0.441 ถึง 0.883 ซึ่งมากกว่าค่าที่กำหนด คือ 0.190 และค่าปัจจัยการขยายตัวของความแปรปรวน (Variance Inflation Factor VIF) มีค่าระหว่าง 1.132 ถึง 2.270 ซึ่งมีค่าไม่เกินตามที่กำหนด คือ 5.30 แสดงว่า การกล่อมเกลาทางการเมืองและหลักสังคหวัตถุ 4 ทุกด้านเป็นไปตามเกณฑ์ที่กำหนดไม่เกิดปัญหา Multicollinearity จึงสามารถวิเคราะห์การถดถอยพหุคูณแบบขั้นตอน (Stepwise Multiple Regression Analysis) ดังนี้

ผลการทดสอบสมมติฐาน

ผลการทดสอบสมมติฐานสำหรับการวิจัยครั้งนี้ มีรายละเอียดดังนี้

สมมติฐานที่ 1 การกล่อมเกลาทางการเมือง ประกอบด้วย 1) ด้านครอบครัว 2) ด้านกลุ่มเพื่อน 3) ด้านสถาบันการศึกษา 4) กลุ่มทางสังคม อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

ในการทดสอบสมมติฐานเกี่ยวกับสมมติฐานที่ 1 เกี่ยวกับการกล่อมเกลาทางการเมืองใช้การวิเคราะห์การถดถอยพหุคูณ (Multiple Regression Analysis) แบบ Stepwise โดยมีตัวแปรที่ศึกษา คือ 1) ด้านครอบครัว 2) ด้านกลุ่มเพื่อน 3) ด้านสถาบันการศึกษา 4) กลุ่มทางสังคม

การวิเคราะห์การถดถอยพหุคูณเพื่อพิจารณาว่า เมื่อทำการวิเคราะห์การถดถอยพหุคูณ (Multiple Regression Analysis) แบบ Stepwise โดยใช้ตัวแปรอิสระจำนวน 4 ตัวแปร แล้วสามารถอธิบายความแปรปรวนของตัวแปรตามได้ร้อยละเท่าใด และมีตัวแปรใดบ้างที่ส่งผลต่อตัวแปรตามอย่างมีนัยสำคัญ รายละเอียดผลการวิเคราะห์ ดังแสดงในตารางที่ 10.9 ดังนี้

ตารางที่ 10.9 การกล่อมเกลาทางการเมือง ประกอบด้วย 1) ด้านครอบครัว 2) ด้านกลุ่มเพื่อน 3) ด้านสถาบันการศึกษา 4) ด้านกลุ่มทางสังคม อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

การกล่อมเกลาทางการเมือง	B	Std.Error	Beta	t	Sig.
ค่าคงที่ (Constant)	.981	.126		7.770**	.000
ด้านกลุ่มเพื่อน _{SumB}	.303	.024	.477	12.740**	.000
ด้านกลุ่มทางสังคม _{SumD}	.205	.018	.417	11.351**	.000
ด้านสถาบันการศึกษา _{SumC}	.205	.024	.324	8.669**	.000

Multiple R = .686^c R Square = .471 Adjusted R Square = .467 Std. Error = .291

** มีนัยสำคัญทางสถิติที่ระดับ 0.01

จากตารางที่ 10.9 พบว่า การกล่อมเกลาทางการเมืองส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยาส่งผลอยู่ 3 ด้านอย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 โดยเรียงลำดับตามสมการดังนี้ คือ ด้านกลุ่มเพื่อน ด้านกลุ่มทางสังคม ด้านครอบครัว พบว่า ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

ผลการวิเคราะห์ค่าสัมประสิทธิ์สหสัมพันธ์พหุคูณ (Multiple R) เท่ากับ .686 ค่าสัมประสิทธิ์ การตัดสินใจ (R Square) เท่ากับ .471 สัมประสิทธิ์การตัดสินใจที่ปรับแล้ว (Adjusted R Square) เท่ากับ .467 และค่าความคลาดเคลื่อนมาตรฐานในการตัดสินใจ (Std. Error) เท่ากับ .291

แสดงว่าการกล่อมเกลாதางการเมือง สามารถร่วมกันทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 46.7 และเมื่อพิจารณาเป็น รายด้าน พบว่า ด้านกลุ่มเพื่อน สามารถทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 47.7 ด้านกลุ่มทางสังคม สามารถทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 41.7 และด้านสถาบันการศึกษา สามารถทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 32.4 ตามลำดับ ซึ่งสามารถเขียนสมการได้ดังนี้

สมการพยากรณ์ในรูปแบบคะแนนดิบ

$$\hat{Y} = .981 + .303(\text{SumB}) + .205(\text{SumD}) + .205(\text{SumC})$$

สมการพยากรณ์ในรูปแบบคะแนนมาตรฐาน

$$Z_y = .477Z(\text{SumB}) + .417Z(\text{SumD}) + .324Z(\text{SumC})$$

สมมติฐานที่ 2 หลักสังเกต 4 ประกอบด้วย 1) ทาน โอบอ้อมอารี 2) ปิยวาจา วจีไพเราะ 3) อุตถจริยา สงเคราะห์ประชาชน 4) สมานัตตตา วางตนพอดี อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

ในการทดสอบสมมติฐานเกี่ยวกับสมมติฐานที่ 2 หลักสังเกต 4 ประกอบด้วย 1) ทาน โอบอ้อมอารี 2) ปิยวาจา วจีไพเราะ 3) อุตถจริยา สงเคราะห์ประชาชน 4) สมานัตตตา วางตนพอดี อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ใช้การวิเคราะห์การถดถอยพหุคูณ (Multiple Regression Analysis) แบบ Stepwise โดยมีตัวแปรที่ศึกษา คือ 1) ทาน โอบอ้อมอารี 2) ปิยวาจา วจีไพเราะ 3) อุตถจริยา สงเคราะห์ประชาชน 4) สมานัตตตา วางตนพอดี

การวิเคราะห์การถดถอยพหุคูณเพื่อพิจารณาว่า เมื่อทำการวิเคราะห์การถดถอยพหุคูณ (Multiple Regression Analysis) แบบ Stepwise โดยใช้ตัวแปรอิสระจำนวน 4 ตัวแปร แล้วสามารถอธิบายความแปรปรวนของตัวแปรตามได้ร้อยละเท่าใด และมีตัวแปรใดบ้างที่ส่งผลต่อตัวแปรตามอย่างมีนัยสำคัญ รายละเอียดผลการวิเคราะห์ ดังแสดงในตารางที่ 10.10 ดังนี้

ตารางที่ 10.10 หลักสัจพจน์ 4 ประกอบด้วย 1) ทาน โอบอ้อมอารี 2) ปิยวาจา วจีไพเราะ 3) อุตถจริยา สงเคราะห์ประชาชน 4) สมานัตตตา วางตนพอดี อย่างน้อย 1 ด้าน ที่ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมือง ในจังหวัด พระนครศรีอยุธยา

การกล่อมเกลางานการเมือง	B	Std.Error	Beta	t	Sig.
ค่าคงที่ (Constant)	2.000	.142		14.094**	.000
ปิยวาจา วจีไพเราะ _{SumBB}	.156	.028	.271	5.527**	.000
อุตถจริยา สงเคราะห์ประชาชน _{SumCC}	.199	.041	.238	4.861**	.000

Multiple R = .421^B R Square = .178 Adjusted R Square = .173 Std. Error = .362
 ** มีนัยสำคัญทางสถิติที่ระดับ 0.01

จากตารางที่ 10.10 พบว่า หลักสัจพจน์ 4 ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยาส่งผลอยู่ 2 ด้าน อย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 โดยเรียงลำดับตามสมการดังนี้ คือ ปิยวาจา วจีไพเราะ และอุตถจริยา สงเคราะห์ประชาชน พบว่า ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา

ผลการวิเคราะห์ค่าสัมประสิทธิ์สหสัมพันธ์พหุคูณ (Multiple R) เท่ากับ .421 ค่าสัมประสิทธิ์ การตัดสินใจ (R Square) เท่ากับ .178 สัมประสิทธิ์การตัดสินใจที่ปรับแล้ว (Adjusted R Square) เท่ากับ .173 และค่าความคลาดเคลื่อนมาตรฐานในการตัดสินใจ (Std. Error) เท่ากับ .362

แสดงว่า หลักสัจพจน์ 4 สามารถร่วมกันทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 17.3 และเมื่อพิจารณาเป็นรายด้าน พบว่า ปิยวาจา วจีไพเราะ สามารถทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 27.1 และอุตถจริยา สงเคราะห์ประชาชน สามารถทำนายความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา ได้ร้อยละ 23.8 ตามลำดับ ซึ่งสามารถเขียนสมการได้ดังนี้

สมการพยากรณ์ในรูปแบบคะแนนดิบ

$$\hat{Y} = 2.000 + .156(\text{SumBB}) + .199(\text{SumCC})$$

สมการพยากรณ์ในรูปแบบคะแนนมาตรฐาน

$$Z_y = .271Z_{(\text{SumBB})} + .238Z_{(\text{SumCC})}$$

สรุปได้ว่า

1) การกล่อมเกลາทางการเมืองส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยาส่งผลอยู่ 3 ด้าน อย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 โดยเรียงลำดับตามสมการดังนี้ แสดงว่าได้ร้อยละ 46.7 และเมื่อพิจารณาเป็นรายด้าน พบว่า ด้านกลุ่มเพื่อน ได้ร้อยละ 47.7 ด้านกลุ่มทางสังคม ได้ร้อยละ 41.7 และด้านสถาบันการศึกษา ได้ร้อยละ 32.4 ตามลำดับ

2) หลักสังคหวัตถุ 4 ส่งผลต่อความนิยมทางการเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยาทุกด้าน อย่างมีนัยสำคัญที่ระดับ 0.01 แสดงว่าได้ร้อยละ 17.3 และเมื่อพิจารณาเป็น รายด้าน พบว่า ปิยวาจา วชิไพบระ ได้ร้อยละ 27.1 และอรรถจริยา สงเคราะห์ประชาชน ได้ร้อยละ 23.8 ตามลำดับ

บทบาทของสถิติเชิงอนุมานในงานวิจัยทางรัฐประศาสนศาสตร์

สถิติเชิงอนุมานมีบทบาทสำคัญอย่างยิ่งในงานวิจัยทางรัฐประศาสนศาสตร์ เนื่องจากช่วยให้นักวิจัยก้าวข้ามจากการบรรยายข้อมูลเชิงพรรณนา ไปสู่การสร้างข้อสรุปที่สามารถตรวจสอบได้ตามหลักวิทยาศาสตร์ และนำไปใช้ประโยชน์ในเชิงนโยบายและการบริหารงานภาครัฐได้จริง ความสำคัญของสถิติเชิงอนุมานสามารถอธิบายเพิ่มเติมได้ดังนี้

1. การเปรียบเทียบกลุ่ม (t-test และ ANOVA)

- การใช้ t-test ช่วยให้สามารถตรวจสอบความแตกต่างของค่าเฉลี่ยระหว่างสองกลุ่ม เช่น ความแตกต่างด้านความพึงพอใจต่อการบริการสาธารณะระหว่างเพศชายและเพศหญิง หรือระหว่างข้าราชการประจำและพนักงานจ้าง

- ขณะที่ ANOVA ช่วยให้เปรียบเทียบความแตกต่างระหว่างกลุ่มมากกว่าสองกลุ่ม เช่น การเปรียบเทียบความคิดเห็นของบุคลากรที่มีระดับการศึกษาต่างกันต่อการบริหารงานขององค์กรปกครองส่วนท้องถิ่น การใช้สถิติทั้งสองนี้ทำให้นักวิจัยสามารถยืนยันได้ว่าความแตกต่างที่พบเกิดจากปัจจัยที่ศึกษา มิใช่ความบังเอิญ

2. การตรวจสอบความสัมพันธ์ (Correlation Analysis)

- ช่วยอธิบายว่าตัวแปรสองตัวหรือมากกว่านั้นมีความสัมพันธ์กันหรือไม่ และมีทิศทางเชิงบวกหรือเชิงลบ เช่น การตรวจสอบความสัมพันธ์ระหว่างการมีส่วนร่วมของประชาชนกับระดับความเชื่อมั่นต่อรัฐบาล

- การใช้สหสัมพันธ์ยังช่วยให้นักวิจัยเข้าใจถึงโครงสร้างของปัญหาและปัจจัยที่เกี่ยวข้อง ซึ่งมีประโยชน์ต่อการออกแบบนโยบายที่ตอบสนองต่อความต้องการของประชาชนได้ตรงจุด

3. การอธิบายและทำนาย (Regression Analysis)

- การถดถอย (Regression) เป็นสถิติที่ใช้ทำนายและอธิบายอิทธิพลของตัวแปรอิสระต่อการเปลี่ยนแปลงของตัวแปรตาม เช่น ศึกษาว่าการมีส่วนร่วมทางการเมืองรายได้ และระดับการศึกษาของประชาชนมีผลต่อความพึงพอใจต่อการให้บริการสาธารณะเพียงใด

- ข้อดีของการใช้ Regression คือสามารถแยกอิทธิพลของแต่ละตัวแปรออกมาได้ ทำให้ผู้กำหนดนโยบายเข้าใจปัจจัยสำคัญที่ควรให้ความสำคัญเป็นพิเศษ

4. การเชื่อมโยงกับเชิงนโยบาย

ผลลัพธ์จากการวิเคราะห์สถิติเหล่านี้ไม่ได้หยุดอยู่เพียงในระดับการวิจัยเชิงวิชาการ แต่ยังสามารถใช้เป็นหลักฐานเชิงประจักษ์ (Evidence-based) เพื่อสนับสนุนการตัดสินใจเชิงนโยบาย เช่น การจัดสรรงบประมาณ การปรับปรุงกระบวนการบริหาร หรือการออกแบบโครงการที่ตอบสนองต่อความต้องการของประชาชนอย่างแท้จริง

สรุปได้ว่า สถิติเชิงอนุมานจึงมีบทบาทเป็นสะพานเชื่อมระหว่าง ข้อมูลเชิงตัวเลขกับการตัดสินใจเชิงนโยบาย โดย t-test และ ANOVA ทำให้เข้าใจความแตกต่างระหว่างกลุ่ม Correlation ช่วยให้เห็นความสัมพันธ์ระหว่างปัจจัย ส่วน Regression ช่วยอธิบายและทำนายผลลัพธ์ที่เกิดขึ้น ตัวอย่างจากงานวิจัยจริงในสาขารัฐประศาสนศาสตร์ ยืนยันแล้วว่า การใช้สถิติเหล่านี้ช่วยเพิ่มความน่าเชื่อถือของผลวิจัย และทำให้ข้อค้นพบมีคุณค่าเชิงปฏิบัติ สามารถนำไปใช้เพื่อสนับสนุนการบริหารจัดการภาครัฐได้อย่างมีระบบและมีหลักฐานรองรับ

3. การใช้สถิติขั้นสูงในงานวิจัยรัฐประศาสนศาสตร์

นอกจากสถิติเชิงพรรณนาและสถิติเชิงอนุมานแล้ว งานวิจัยทางรัฐประศาสนศาสตร์สมัยใหม่มีแนวโน้มใช้สถิติขั้นสูง (Advanced Statistics) เพื่อรองรับความซับซ้อนของข้อมูลและตรวจสอบความสัมพันธ์เชิงสาเหตุที่ซับซ้อนของตัวแปรทั้งที่สังเกตได้ (Observed Variables) และตัวแปรแฝง (Latent Variables) วิธีการเหล่านี้ช่วยให้ผลการวิจัยมีความน่าเชื่อถือมากขึ้น และสามารถสร้างโมเดลเชิงทฤษฎีที่รองรับการกำหนดนโยบายและการบริหารจัดการภาครัฐได้จริง

1) การวิเคราะห์องค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis: CFA)

การวิเคราะห์องค์ประกอบเชิงยืนยันเป็นสถิติขั้นสูงที่ใช้ตรวจสอบว่า โครงสร้างตัวแปรแฝง (Latent Variables) และตัวชี้วัด (Indicators) ที่นักวิจัยกำหนดขึ้นมีความสอดคล้องกับทฤษฎี กรอบแนวคิด หรือสมมติฐาน ที่วางไว้หรือไม่

จุดประสงค์หลักของ CFA

- ตรวจสอบความตรงเชิงโครงสร้าง (Construct Validity) ของเครื่องมือวิจัย
- ยืนยันว่าตัวชี้วัด ที่สร้างขึ้นสะท้อนมิติของตัวแปรแฝงตามทฤษฎีได้อย่าง

เหมาะสม

- ประเมินคุณภาพของแบบสอบถาม โดยดูค่าความเที่ยงตรงและความเชื่อมั่น

ดัชนีที่นิยมใช้ในการประเมินโมเดล

- Chi-Square (χ^2) ถ้าไม่แตกต่างจาก 0 หรือค่า p-value > .05 แสดงว่าโมเดลเหมาะสม

- GFI (Goodness of Fit Index), AGFI ควรมากกว่า 0.90
- CFI (Comparative Fit Index) ควรมากกว่า 0.90
- RMSEA (Root Mean Square Error of Approximation) ควรน้อยกว่า 0.08

ตัวอย่างงานวิจัย

งานวิจัยของสุภัทรชัย สีสะไบ (2563) เรื่อง การพัฒนาโมเดลประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ได้ประยุกต์ใช้การวิเคราะห์องค์ประกอบเชิงยืนยัน (CFA) เพื่อตรวจสอบความตรงเชิงโครงสร้าง (Construct Validity) ของโมเดลที่พัฒนาขึ้น มีความสอดคล้องกับโมเดลเชิงทฤษฎีที่กำหนดไว้หรือไม่

การตรวจสอบความตรงเชิงโครงสร้าง (Construct Validity)

ในการตรวจสอบความตรงเชิงโครงสร้างของข้อมูลที่ได้ ผู้วิจัยได้ใช้วิธีตรวจสอบความตรงเชิงโครงสร้าง (Construct Validity) โดยการวิเคราะห์สหสัมพันธ์ระหว่างตัวแปรให้ได้เมทริกซ์สหสัมพันธ์สหสัมพันธ์ระหว่างตัวแปรในแต่ละองค์ประกอบ โดยมีวัตถุประสงค์เพื่อตรวจสอบว่าเมทริกซ์สหสัมพันธ์สหสัมพันธ์แตกต่างจากศูนย์หรือไม่ ถ้าสหสัมพันธ์สหสัมพันธ์ในเมทริกซ์ใดไม่มีความสัมพันธ์กัน หรือมีความสัมพันธ์กันน้อย แสดงว่าเมทริกซ์นั้นไม่มีองค์ประกอบร่วมกันและไม่จำเป็นในการนำเมทริกซ์สหสัมพันธ์สหสัมพันธ์ไปวิเคราะห์องค์ประกอบ สำหรับค่าสถิติที่ใช้ในการทดสอบสมมติฐานคือค่าสถิติ Bartlett's Test of Sphericity และค่าดัชนีไกเซอร์ เมเยอร์-ออลคิน (Kaiser-

Meyer-Olkin Measure of Sampling Adequacy = KMO) ซึ่งค่า KMO ควรจะมีค่าเข้าใกล้หนึ่ง ถ้ามีค่าน้อยแสดงว่าความสัมพันธ์ระหว่างตัวแปรมีน้อย และไม่เหมาะที่จะวิเคราะห์องค์ประกอบ รายละเอียดเกณฑ์ค่าดัชนี KMO มีดังนี้

ค่าดัชนี Kaiser-Meyer-Olkin (KMO)	ระดับความเหมาะสม
KMO > .90	ดีมาก
.80 < KMO < .89	ดี
.70 < KMO < .79	ปานกลาง
.60 < KMO < .69	น้อย
.50 < KMO < .59	น้อยมาก
KMO < .50	ไม่เหมาะสม

ในการตรวจสอบความตรงเชิงโครงสร้างและการวิเคราะห์เพื่อตรวจสอบความตรงของโมเดลแบบมีตัวแปรส่งผ่านด้วยโปรแกรม LISREL จำเป็นต้องมีการเตรียมเมทริกซ์สัมประสิทธิ์สหสัมพันธ์ระหว่างตัวแปรของแต่ละองค์ประกอบ และในการแปลความหมายของค่าสัมประสิทธิ์สหสัมพันธ์สำหรับการวิจัยครั้งนี้ใช้การแปลความหมายของขนาดความสัมพันธ์ ดังต่อไปนี้

ขนาดความสัมพันธ์	ความหมาย
0.0 – 0.3	มีความสัมพันธ์กันต่ำมาก
0.3 – 0.5	มีความสัมพันธ์กันต่ำ
0.5 – 0.7	มีความสัมพันธ์กันปานกลาง
0.7 – 0.9	มีความสัมพันธ์กันสูง
0.9 – 1.0	มีความสัมพันธ์กันสูงมาก

เมื่อได้เมทริกซ์สัมประสิทธิ์สหสัมพันธ์ระหว่างตัวแปรแต่ละองค์ประกอบ จากนั้นผู้วิจัยได้นำมาวิเคราะห์เพื่อเป็นการตรวจสอบองค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis: CFA) ด้วยโปรแกรมสถิติขั้นสูง โดยใช้แบบสอบถามที่ได้จากการเก็บรวบรวมข้อมูลจากกลุ่มตัวอย่างจำนวน 60 คน แสดงผลการวิเคราะห์ได้ดังต่อไปนี้

ความตรงเชิงโครงสร้างขององค์ประกอบหลักการบริหารจัดการแบบมีส่วนร่วม

ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรโดยใช้ค่าสหสัมพันธ์แบบเพียร์สันพบว่า ตัวแปรที่บ่งชี้องค์ประกอบหลักการบริหารจัดการแบบมีส่วนร่วม มีค่าสัมประสิทธิ์สหสัมพันธ์ตั้งแต่ .826 ถึง .958 อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกคู่ ลักษณะความสัมพันธ์ระหว่างตัวแปรเป็นความสัมพันธ์ทางบวกตั้งแต่ขนาดปานกลางถึงขนาดสูง ตัวแปรที่มีความสัมพันธ์ทางบวกสูงที่สุดคือ ร่วมวางแผน (Management3) กับร่วมดำเนินการ (Management4) และตัวแปรที่มีความสัมพันธ์ทางบวกต่ำที่สุด คือ ร่วมติดตามประเมินผล (Management5) กับร่วมรับผลประโยชน์ (Management6) เมื่อพิจารณาค่า Bartlett's Test of Sphericity มีค่าเท่ากับ 631.457 ($p = .000$) แสดงว่า เมทริกซ์สหสัมพันธ์ระหว่างตัวแปรแตกต่างจากเมทริกซ์เอกลักษณ์อย่างมีนัยสำคัญ และค่าดัชนีไกเซอร์-ไมเยอร์-ออลกิน (Kaiser-Meyer-Olkin Measure of Sampling Adequacy) มีค่าเท่ากับ .925 แสดงว่าตัวแปรสังเกตได้ของข้อมูลมีความสัมพันธ์กันมากพอที่จะนำไปวิเคราะห์องค์ประกอบได้ ดังแสดงในตารางที่ 10.11

ตารางที่ 10.11 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน ระหว่างตัวแปรในองค์ประกอบหลักการบริหารจัดการแบบมีส่วนร่วม

ตัวแปร	Management1	Management2	Management3	Management4	Management5	Management6
Management1	1.000					
Management2	.942**	1.000				
Management3	.939**	.944**	1.000			
Management4	.930**	.934**	.958**	1.000		
Management5	.880**	.876**	.947**	.923**	1.000	
Management6	.874**	.875**	.871**	.876**	.826**	1.000
MEAN	2.825	2.819	2.750	2.838	2.833	2.917
SD	1.078	1.128	1.134	1.114	1.213	1.187

Kaiser-Meyer-Olkin Measure of Sampling Adequacy = .925

Bartlett's Test of Sphericity = 631.457, $df = 15$, $p = .000$

หมายเหตุ: ** $p < .01$; $n = 60$; ระดับการให้คะแนน 5 ระดับ

ผลการวิเคราะห์องค์ประกอบยืนยันตามโมเดลการวัดองค์ประกอบหลักการบริหารจัดการพบว่า โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ดี พิจารณาได้

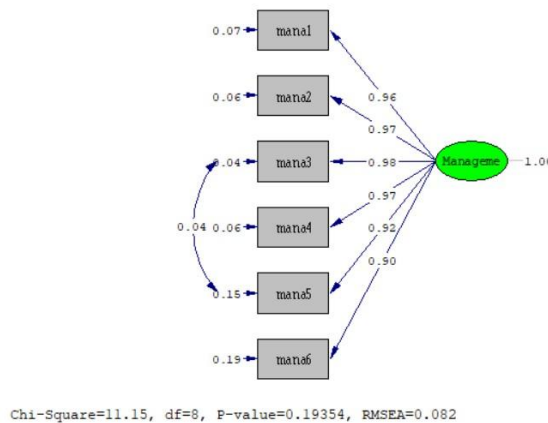
จากค่าไค-สแควร์ ($\chi^2 = 11.15$, $df = 8$, $p = 0.193$) ซึ่งแตกต่างจากศูนย์อย่างไม่มีนัยสำคัญ ค่าดัชนีวัดระดับความกลมกลืน (GFI) มีค่าเท่ากับ 0.94 ค่าดัชนีวัดระดับความกลมกลืนที่ปรับแก้แล้ว (AGFI) มีค่าเท่ากับ 0.84 และค่ารากที่สองของค่าเฉลี่ยความคลาดเคลื่อนกำลังสองของการประมาณค่า (RMSEA) มีค่าเท่ากับ 0.082 แสดงว่า โมเดลมีความสอดคล้องกับข้อมูลเชิงประจักษ์

ค่าน้ำหนักองค์ประกอบของตัวแปรทั้งหมดมีค่าเป็นบวก มีขนาดตั้งแต่ 0.90 ถึง 0.98 และมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกตัว ตัวแปรที่มีน้ำหนักความสำคัญมากที่สุด คือ ร่วมวางแผน (Management3) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.98 และมีความแปรผันร่วมกับหลักการบริหารจัดการแบบมีส่วนร่วม ร้อยละ 96 รองลงมาคือ ร่วมตัดสินใจ (Management2) และร่วมดำเนินการ (Management4) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.97 และมีความแปรผันร่วมกับหลักการบริหารจัดการ ร้อยละ 94 และตัวแปรที่มีน้ำหนักความสำคัญน้อยที่สุด คือ ร่วมรับผลประโยชน์ (Management6) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.90 และมีความแปรผันร่วมกับหลักการบริหารจัดการ ร้อยละ 81 แสดงให้เห็นว่าตัวแปรเหล่านี้เป็นตัวแปรที่สำคัญขององค์ประกอบหลักการบริหารจัดการแบบมีส่วนร่วม ดังแสดงในตารางที่ 10.12 และแผนภาพที่ 10.1

ตารางที่ 10.12 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดองค์ประกอบหลักการบริหารจัดการแบบมีส่วนร่วม

ตัวแปร	น้ำหนักองค์ประกอบ		t	R ²	สปส.คะแนน องค์ประกอบ
	beta	B(SE)			
Management1	0.96	1.04(0.10)	10.09**	0.93	0.17
Management2	0.97	1.09(0.11)	10.18**	0.94	0.19
Management3	0.98	1.11(0.11)	10.40**	0.96	0.29
Management4	0.97	1.08(0.11)	10.26**	0.94	0.22
Management5	0.92	1.12(0.12)	9.35**	0.85	-0.01
Management6	0.90	1.07(0.12)	8.93**	0.81	0.05
$\chi^2 = 13.39$ $df = 8$ $p = 0.099$ $GFI = 0.94$ $AGFI = 0.84$ $RMSEA = 0.11$					

หมายเหตุ: ** $p < .01$



แผนภาพที่ 10.1 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัย
หลักการบริหารจัดการแบบมีส่วนร่วม

ความตรงเชิงโครงสร้างขององค์ประกอบปัจจัยแห่งความสำเร็จ

ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรโดยใช้ค่าสัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน พบว่า ตัวแปรที่บ่งชี้องค์ประกอบการมีส่วนร่วมของประชาชน มีค่าสัมประสิทธิ์สหสัมพันธ์ตั้งแต่ 0.750 ถึง 0.951 อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกคู่ ลักษณะความสัมพันธ์ระหว่างตัวแปรเป็นความสัมพันธ์ทางบวกตั้งแต่ขนาดต่ำถึงขนาดสูง ตัวแปรที่มีความสัมพันธ์ทางบวกสูงที่สุดคือ ทักษะ/ความสามารถ (Succession2) กับ เทคโนโลยี/นวัตกรรม (Succession3) และตัวแปรที่มีความสัมพันธ์ทางบวกต่ำที่สุดคือ ความรู้ (Succession1) กับ ทุน (Succession5) เมื่อพิจารณาค่า Bartlett's Test of Sphericity มีค่าเท่ากับ 532.455 ($p = .000$) แสดงว่าเมทริกซ์สหสัมพันธ์ระหว่างตัวแปรแตกต่างจากเมทริกซ์เอกลักษณ์อย่างมีนัยสำคัญ และค่าดัชนีไกเซอร์ - ไมเยอร์ - ออลคิน (Kaiser-Meyer-Olkin Measure of Sampling Adequacy) มีค่าเท่ากับ .888 แสดงว่าตัวแปรสังเกตได้ของข้อมูลมีความสัมพันธ์กันมากพอที่จะนำไปวิเคราะห์องค์ประกอบได้ ดังแสดงในตารางที่ 10.13

ตารางที่ 10.13 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน
ระหว่างตัวแปรในองค์ประกอบปัจจัยแห่งความสำเร็จ

ตัวแปร	Succession1	Succession2	Succession3	Succession4	Succession5	Succession6
Succession1	1.000					
Succession2	.928**	1.000				
Succession3	.917**	.951**	1.000			
Succession4	.855**	.899**	.911**	1.000		
Succession5	.750**	.806**	.859**	.800**	1.000	
Succession6	.775**	.818**	.842**	.840**	.891**	1.000
MEAN	3.283	3.342	3.311	3.411	3.008	3.121
SD	1.109	1.039	1.076	1.089	1.083	1.117

Kaiser-Meyer-Olkin Measure of Sampling Adequacy = .888

Bartlett's Test of Sphericity = 532.455, df = 15, p = .000

หมายเหตุ: **p < .01; n = 60; ระดับการให้คะแนน 5 ระดับ

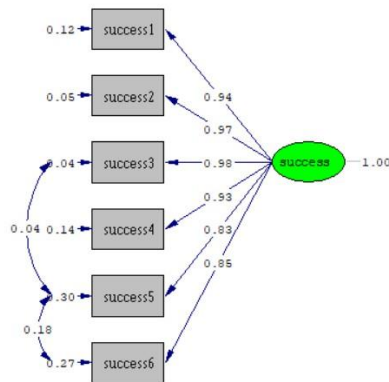
ผลการวิเคราะห์องค์ประกอบยืนยันตามโมเดลการวัดองค์ประกอบปัจจัยแห่งความสำเร็จ พบว่า โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ดี พิจารณาได้จากค่า ไค-สแควร์ ($\chi^2 = 10.02$, df = 7, p = .187) ซึ่งแตกต่างจากศูนย์อย่างไม่มีนัยสำคัญ ค่าดัชนีวัดระดับความกลมกลืน (GFI) มีค่าเท่ากับ 0.95 ค่าดัชนีวัดระดับความกลมกลืนที่ปรับแก้แล้ว (AGFI) มีค่าเท่ากับ 0.84 และค่ารากที่สองของค่าเฉลี่ยความคลาดเคลื่อนกำลังสองของการประมาณค่า (RMSEA) มีค่าเท่ากับ 0.086 แสดงว่า โมเดลมีความสอดคล้องกับข้อมูลเชิงประจักษ์

ค่าน้ำหนักองค์ประกอบของตัวแปรทั้งหมดมีค่าเป็นบวก มีขนาดตั้งแต่ 0.83 ถึง 0.98 และมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกตัว ตัวแปรที่มีน้ำหนักความสำคัญมากที่สุดคือเทคโนโลยี/นวัตกรรม (Succession3) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.98 และมีความแปรผันร่วมกับปัจจัยแห่งความสำเร็จร้อยละ 96 รองลงมา คือ ทักษะ/ความสามารถ (Succession2) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.97 และมีความแปรผันร่วมกับปัจจัยแห่งความสำเร็จ ร้อยละ 95 และตัวแปรที่มีน้ำหนักความสำคัญน้อยที่สุด คือ พุน (Succession5) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.83 และมีความแปรผันร่วมกับปัจจัยแห่งความสำเร็จ ร้อยละ 70 แสดงให้เห็นว่า ตัวแปรเหล่านี้เป็นตัวแปรที่สำคัญขององค์ประกอบปัจจัยแห่งความสำเร็จ ดังแสดงในตารางที่ 10.14 และแผนภาพที่ 10.2

ตารางที่ 10.14 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดปัจจัยแห่งความสำเร็จ

ตัวแปร	น้ำหนักองค์ประกอบ		t	R ²	สปส.คะแนนองค์ประกอบ
	beta	B(SE)			
Succession1	0.94	1.04(0.110)	9.65**	0.88	0.12
Succession2	0.97	1.01(0.098)	10.27**	0.95	0.29
Succession3	0.98	1.05(0.100)	10.38**	0.96	0.41
Succession4	0.93	1.01(0.110)	9.44**	0.86	0.10
Succession5	0.83	0.90(0.110)	7.89**	0.70	-0.07
Succession6	0.85	0.95(0.120)	8.20**	0.73	0.09
$\chi^2 = 10.02$ df = 7 p = 0.187 GFI = 0.95 AGFI = 0.84 RMSEA = 0.086					

หมายเหตุ: **p < .01



Chi-Square=10.02, df=7, P-value=0.18741, RMSEA=0.086

แผนภาพที่ 10.2 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยแห่งความสำเร็จ

ความตรงเชิงโครงสร้างขององค์ประกอบธรรมะเพื่อการปฏิบัติ

ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรโดยใช้ค่าสัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน พบว่า ตัวแปรที่บ่งชี้องค์ประกอบพฤติกรรมบริหารจัดการ มีค่าสัมประสิทธิ์สหสัมพันธ์ตั้งแต่ .858 ถึง .951 อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกคู่ ลักษณะความสัมพันธ์ระหว่างตัวแปรเป็นความสัมพันธ์ทางบวกตั้งแต่ขนาดต่ำถึงขนาดสูง ตัวแปรที่มีความสัมพันธ์ทางบวกสูงที่สุดคือ พากเพียรทำ (วิริยะ) กับนำจิตฝึกฝน (จิตตะ) และตัวแปรที่มี

ความสัมพันธ์ทางบวกต่ำที่สุดคือ พากเพียรทำ (วิริยะ) กับใช้ปัญญาทบทวน (วิมังสา) เมื่อพิจารณาค่า Bartlett's Test of Sphericity มีค่าเท่ากับ 376.375 ($p = .000$) แสดงว่าเมทริกซ์สหสัมพันธ์ระหว่างตัวแปรแตกต่างจากเมทริกซ์เอกลักษณ์อย่างมีนัยสำคัญ และค่าดัชนีไกเซอร์ - ไมเยอร์ - ออลคิน (Kaiser-Meyer-Olkin Measure of Sampling Adequacy) มีค่าเท่ากับ .848 แสดงว่า ตัวแปรสังเกตได้ของข้อมูลมีความสัมพันธ์กันมากพอที่จะนำไปวิเคราะห์องค์ประกอบได้ ดังแสดงในตารางที่ 10.15

ตารางที่ 10.15 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน ระหว่างตัวแปรในองค์ประกอบธรรมชาติเพื่อการปฏิบัติ

ตัวแปร	Dhamma1	Dhamma2	Dhamma3	Dhamma4
Dhamma1	1.000			
Dhamma2	.950**	1.000		
Dhamma3	.944**	.951**	1.000	
Dhamma4	.884**	.858**	.903**	1.000
MEAN	3.603	3.504	3.423	3.317
SD	1.129	1.119	1.108	1.069

Kaiser-Meyer-Olkin Measure of Sampling Adequacy = .848
Bartlett's Test of Sphericity = 376.375, $df = 6$, $p = .000$

หมายเหตุ: ** $p < .01$; $n = 60$; ระดับการให้คะแนน 5 ระดับ

ผลการวิเคราะห์องค์ประกอบยืนยันตามโมเดลการวัดองค์ประกอบธรรมชาติเพื่อการปฏิบัติ พบว่า โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ดี พิจารณาได้จากค่าไค-สแควร์ ($\chi^2 = 0.64$, $df = 1$, $p = 0.424$) ซึ่งแตกต่างจากศูนย์อย่างไม่มีนัยสำคัญ ค่าดัชนีวัดระดับความกลมกลืน (GFI) มีค่าเท่ากับ 0.99 ค่าดัชนีวัดระดับความกลมกลืนที่ปรับแก้แล้ว (AGFI) มีค่าเท่ากับ 0.95 และค่ารากที่สองของค่าเฉลี่ยความคลาดเคลื่อนกำลังสองของการประมาณค่า (RMSEA) มีค่าเท่ากับ 0.000 แสดงว่า โมเดลมีความสอดคล้องกับข้อมูลเชิงประจักษ์

ค่าน้ำหนักองค์ประกอบของตัวแปรทั้งหมดมีค่าเป็นบวก มีขนาดตั้งแต่ 0.92 ถึง 0.98 และมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกตัว ตัวแปรที่มีน้ำหนักความสำคัญมากที่สุด คือ พากเพียรทำ (วิริยะ) (Dhamma2) และนำจิตฝึกฝน (จิตตะ) (Dhamma3) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.98 เท่ากันทั้งสองตัวแปร และมีความแปรผันร่วมกับปัจจัยธรรมชาติเพื่อการปฏิบัติร้อยละ 96 และตัวแปรที่มีน้ำหนักความสำคัญน้อยที่สุด คือ ใช้ปัญญา

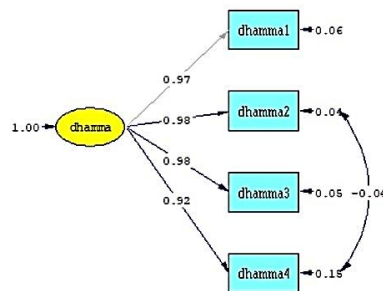
ทบทวน (วิมังสา) (Dhamma4) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.92 และมีความแปรผันร่วมกับธรรมะเพื่อการปฏิบัติ ร้อยละ 84 แสดงให้เห็นว่า ตัวแปรเหล่านี้เป็นตัวแปรที่สำคัญขององค์ประกอบธรรมะเพื่อการปฏิบัติ ดังแสดงในตารางที่ 10.16 และแผนภาพที่ 10.3

ตารางที่ 10.16 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดองค์ประกอบธรรมะเพื่อการปฏิบัติ

ตัวแปร	น้ำหนักองค์ประกอบ		t	R ²	สปส.คะแนนองค์ประกอบ
	beta	B(SE)			
Dhamma1	0.97	1.09		0.94	0.16
Dhamma2	0.98	1.09(0.05)	22.35**	0.96	0.38
Dhamma3	0.98	1.08(0.05)	22.12**	0.96	0.21
Dhamma4	0.92	0.99(0.07)	15.15**	0.85	0.17

$\chi^2 = 0.64$ df = 1 p = 0.425 GFI = 0.99 AGFI = 0.95 RMSEA = 0.00

หมายเหตุ: **p < .01



Chi-Square=0.64, df=1, P-value=0.42456, RMSEA=0.000

แผนภาพที่ 10.3 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยธรรมะเพื่อการปฏิบัติ

ความตรงเชิงโครงสร้างขององค์ประกอบประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร

ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรโดยใช้ค่าสัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน พบว่า ตัวแปรที่บ่งชี้องค์ประกอบประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์ฯ มีค่าสัมประสิทธิ์สหสัมพันธ์ตั้งแต่ .931 ถึง .977 อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกคู่ ลักษณะความสัมพันธ์ระหว่างตัวแปรเป็นความสัมพันธ์ทางบวกตั้งแต่ขนาดปานกลางถึง

ขนาดสูง ตัวแปรที่มีความสัมพันธ์กันสูงที่สุดคือดำรงอยู่ไม่เสื่อมคลาย (Effective3) กับยอมรับได้ในสังคม (Effective4) และตัวแปรที่มีความสัมพันธ์กันต่ำที่สุดคือ พัฒนาไม่หยุดยั้ง (Effective1) กับยอมรับได้ในสังคม (Effective4) เมื่อพิจารณาค่า Bartlett's Test of Sphericity มีค่าเท่ากับ 472.155 ($p = .000$) แสดงว่า เมทริกซ์สหสัมพันธ์ระหว่างตัวแปรแตกต่างจากเมทริกซ์เอกลักษณ์อย่างมีนัยสำคัญ และค่าดัชนี ไกเซอร์ - ไมเยอร์ - ออลคิน (Kaiser-Meyer-Olkin Measure of Sampling Adequacy) มีค่าเท่ากับ .848 แสดงว่า ตัวแปรสังเกตได้ของข้อมูลมีความสัมพันธ์กันมากพอที่จะนำไปวิเคราะห์องค์ประกอบได้ ดังแสดงในตารางที่ 10.17

ตารางที่ 10.17 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน สัมประสิทธิ์สหสัมพันธ์แบบเพียร์สัน ระหว่างตัวแปรในองค์ประกอบประสิทธิผลพฤติกรรมวิธีการบริหารจัดการ ศูนย์ฯ

ตัวแปร	Effective1	Effective2	Effective3	Effective4
Effective1	1.000			
Effective2	.958**	1.000		
Effective3	.935**	.962**	1.000	
Effective4	.931**	.957**	.977**	1.000
MEAN	3.360	3.328	3.373	3.275
SD	1.064	1.054	1.108	1.125

Kaiser-Meyer-Olkin Measure of Sampling Adequacy = .848

Bartlett's Test of Sphericity = 472.155, $df = 6$, $p = .000$

หมายเหตุ: ** $p < .01$; $n = 60$; ระดับการให้คะแนน 5 ระดับ

ผลการวิเคราะห์องค์ประกอบยืนยันตามโมเดลการวัดองค์ประกอบประสิทธิผล พฤติกรรมการบริหารจัดการ พบว่า โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ดี พิจารณาได้จากค่าไค-สแควร์ ($\chi^2 = 0.05$, $df = 1$, $p = 0.821$) ซึ่งแตกต่างจากศูนย์ อย่างไม่มีนัยสำคัญ ค่าดัชนีวัดระดับความกลมกลืน (GFI) มีค่าเท่ากับ 1.000 ค่าดัชนีวัดระดับความกลมกลืนที่ปรับแก้แล้ว (AGFI) มีค่าเท่ากับ 1.000 และค่ารากที่สองของค่าเฉลี่ย ความคลาดเคลื่อนกำลังสองของการประมาณค่า (RMSEA) มีค่าเท่ากับ 0.000 แสดงว่า โมเดลมีความสอดคล้องกับข้อมูลเชิงประจักษ์

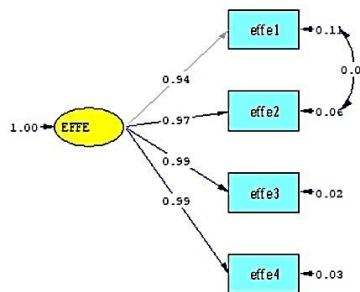
ค่าน้ำหนักองค์ประกอบของตัวแปรทั้งหมดมีค่าเป็นบวก มีขนาดตั้งแต่ 0.94 ถึง 0.99 และมีนัยสำคัญทางสถิติที่ระดับ .01 ทุกตัว ตัวแปรที่มีน้ำหนักความสำคัญมากที่สุด

คือ ดำรงอยู่ไม่เสื่อมคลาย (Effective3) และยอมรับได้ในสังคม (Effective4) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.99 ทั้งสองตัวแปร และมีความแปรผันร่วมกับปัจจัยประสิทธิผลพุทธรวิถีการบริหารจัดการศูนย์ ร้อยละ 98 และตัวแปรที่มีน้ำหนักความสำคัญน้อยที่สุด คือ พัฒนาไม่หยุดยั้ง (Effective1) มีค่าน้ำหนักองค์ประกอบเท่ากับ 0.94 และมีความแปรผันร่วมกับปัจจัยประสิทธิผลพุทธรวิถีการบริหารจัดการศูนย์ ร้อยละ 89 แสดงให้เห็นว่าตัวแปรเหล่านี้เป็นตัวแปรที่สำคัญขององค์ประกอบประสิทธิผลพุทธรวิถีการบริหารจัดการศูนย์ ดังแสดงในตารางที่ 10.18 และแผนภาพที่ 10.4

ตารางที่ 10.18 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันของโมเดลการวัดองค์ประกอบประสิทธิผลพุทธรวิถีการบริหารจัดการ

ตัวแปร	น้ำหนักองค์ประกอบ		t	R ²	สปส.คะแนนองค์ประกอบ
	beta	B(SE)			
Effective1	0.94	1.02		0.89	0.03
Effective2	0.97	1.04(0.041)	25.26**	0.94	0.13
Effective3	0.99	1.11(0.055)	20.26**	0.98	0.46
Effective4	0.99	1.12(0.057)	19.52**	0.98	0.28
$\chi^2 = 0.05$ df = 1 p = 0.821 GFI = 1.000 AGFI = 1.000 RMSEA = 0.000					

หมายเหตุ: **p < .01



Chi-Square=0.05, df=1, P-value=0.82134, RMSEA=0.000

แผนภาพที่ 10.4 ผลการวิเคราะห์องค์ประกอบเชิงยืนยันโมเดลการวัดปัจจัยประสิทธิผลพุทธรวิถีการบริหารจัดการศูนย์

สรุปได้ว่า การวิเคราะห์องค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis: CFA) เป็นเครื่องมือสำคัญที่ใช้ตรวจสอบความตรงเชิงโครงสร้างของเครื่องมือวิจัย โดยเฉพาะในสาขารัฐประศาสนศาสตร์ซึ่งต้องการเครื่องมือที่สามารถสะท้อนมิติของ

การบริหารและนโยบายได้อย่างถูกต้อง งานวิจัยของสุภัทรชัย สีสะใบ (2563) แสดงให้เห็นอย่างชัดเจนว่า การใช้ CFA สามารถยืนยันได้ว่าตัวชี้วัดด้านการบริหารจัดการ เช่น ความรวดเร็ว ความโปร่งใส ความเป็นธรรม รวมถึงมิติของการบริหารจัดการแบบมีส่วนร่วม ปัจจัยแห่งความสำเร็จ ธรรมะเพื่อการปฏิบัติ และประสิทธิผลเชิงพุทธวิธี มีความสอดคล้องกับโมเดลเชิงทฤษฎีที่พัฒนาขึ้น ค่าดัชนีความสอดคล้องของโมเดล (Goodness of Fit Indices) ไม่ว่าจะเป็นค่า KMO, Bartlett's Test, Chi-square, GFI, AGFI และ RMSEA อยู่ในเกณฑ์ที่ยอมรับได้ แสดงให้เห็นถึงความเหมาะสมของโมเดลเชิงทฤษฎีกับข้อมูลเชิงประจักษ์ นอกจากนี้ ค่าน้ำหนักองค์ประกอบ (Factor Loadings) ของตัวแปรทั้งหมดมีค่าตั้งแต่ปานกลางจนถึงสูงมาก (0.83–0.99) และมีนัยสำคัญทางสถิติทุกตัว สะท้อนว่าตัวแปรที่เลือกใช้เป็นตัวบ่งชี้ที่มีคุณภาพสูง

ดังนั้น การใช้ CFA ไม่เพียงช่วยยืนยันความน่าเชื่อถือของเครื่องมือวิจัย แต่ยังเป็นฐานที่มั่นคงสำหรับการพัฒนาโมเดลเชิงทฤษฎีและการประยุกต์ใช้จริงในเชิงการบริหารจัดการและนโยบายสาธารณะ โดยเฉพาะในงานวิจัยรัฐประศาสนศาสตร์ที่มุ่งเน้นการพัฒนา กลไกบริหารงานภาครัฐและศูนย์เรียนรู้ชุมชนให้มีประสิทธิผลและยั่งยืนตามแนวทางเชิงพุทธ

2) การวิเคราะห์สมการโครงสร้าง (Structural Equation Modeling: SEM)

การวิเคราะห์สมการโครงสร้าง (SEM) เป็นสถิติขั้นสูงที่ได้รับความนิยมอย่างมาก ในงานวิจัยด้านรัฐประศาสนศาสตร์ เนื่องจากสามารถตรวจสอบความสัมพันธ์เชิงซับซ้อนระหว่างตัวแปรทั้งที่สังเกตได้ (Observed Variables) และตัวแปรแฝง (Latent Variables) ได้พร้อมกัน โดย SEM เป็นการผสมผสานระหว่าง การวิเคราะห์องค์ประกอบ (Factor Analysis) และการวิเคราะห์การถดถอย (Regression Analysis) ทำให้สามารถทดสอบทั้ง โครงสร้างของโมเดลการวัด (Measurement Model) และโมเดลเชิงโครงสร้าง (Structural Model) ได้ในคราวเดียว

จุดประสงค์หลักของ SEM

- ตรวจสอบความสอดคล้องของโมเดลทฤษฎีกับข้อมูลเชิงประจักษ์
- วิเคราะห์ความสัมพันธ์เชิงสาเหตุ (Causal Relationships) ระหว่างตัวแปร
- พัฒนาและยืนยันโมเดลที่สามารถใช้ทำนายพฤติกรรมหรือปรากฏการณ์ทางสังคม

- ประเมินทั้งตัวชี้วัด (Indicators) และความสัมพันธ์เชิงโครงสร้างของตัวแปรแฝง
- ดัชนีที่นิยมใช้ประเมินความสอดคล้องของโมเดล**

- Chi-square/df (χ^2/df) ควรน้อยกว่า 2 หรือ 3

- GFI (Goodness of Fit Index) และ AGFI ควรมากกว่า 0.90
- CFI (Comparative Fit Index) ควรมากกว่า 0.90
- RMSEA (Root Mean Square Error of Approximation) ควรน้อยกว่า 0.08

ตัวอย่างงานวิจัย

งานวิจัยของสุภัทรชัย สีสะไบ (2563) เรื่อง การพัฒนาโมเดลประสิทธิผลพุทธรวิถีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ได้ใช้ SEM เพื่อทดสอบโมเดลที่ผานปัจจัยด้านการบริหารจัดการแบบมีส่วนร่วม ปัจจัยแห่งความสำเร็จ ปัจจัยธรรมะเพื่อการปฏิบัติ และปัจจัยประสิทธิผลการบริหารจัดการ ว่าสอดคล้องกับข้อมูลเชิงประจักษ์หรือไม่

ผลการวิเคราะห์ SEM

การตรวจสอบความสอดคล้องของโมเดลประสิทธิผลพุทธรวิถีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตรกับข้อมูลเชิงประจักษ์

เป็นการตรวจสอบความสอดคล้องของโมเดลกับข้อมูลเชิงประจักษ์ และวิเคราะห์อิทธิพลทางตรงและทางอ้อมระหว่างตัวแปรในโมเดล ด้วยการวิเคราะห์ข้อมูลเชิงปริมาณ ดังนั้น เพื่อให้การนำเสนอผลการวิเคราะห์ข้อมูลมีความชัดเจนมากยิ่งขึ้น ผู้วิจัยได้กำหนดสัญลักษณ์และอักษรย่อที่ใช้แทนตัวแปรต่างๆ ดังต่อไปนี้

1. สัญลักษณ์ที่ใช้ในการวิจัย

1.1 สัญลักษณ์ที่ใช้แทนค่าสถิติ

$\bar{X}$	หมายถึง ค่าเฉลี่ยเลขคณิต (Mean)
S.D.	หมายถึง ค่าส่วนเบี่ยงเบนมาตรฐาน (Coefficient of Variation)
Max	หมายถึง ค่าสูงสุด (Maximum)
Min	หมายถึง ค่าต่ำสุด (Minimum)
Sk	หมายถึง ค่าความเบ้ (Skewness)
Ku	หมายถึง ค่าความโด่ง (Kurtosis)
χ^2	หมายถึง ดัชนีตรวจสอบความกลมกลืนประเภทค่าสถิติไค-สแควร์
df	หมายถึง องศาอิสระ (Degree of freedom)
p	หมายถึง ระดับนัยสำคัญ (Significant)
TE	หมายถึง ขนาดอิทธิพลรวม (Total effect)
ID	หมายถึง ขนาดอิทธิพลทางอ้อม (Indirect effect)

DE	หมายถึง ขนาดอิทธิพลทางตรง (Direct effect)
R	หมายถึง ค่าสัมประสิทธิ์สหสัมพันธ์พหุคูณ
R ²	หมายถึง สัมประสิทธิ์ การทำนาย (Coefficient of determination)
GFI	หมายถึง ดัชนีวัดระดับความกลมกลืน (Goodness of fit index)
AGFI	หมายถึง ดัชนีวัดระดับความกลมกลืนที่ปรับแก้แล้ว (Adjusted goodness of fit index)
RMR	หมายถึง ค่าดัชนีรากของกำลังสองเฉลี่ยของส่วนที่เหลือ (Root mean square residual)
RMSEA	หมายถึง ค่าดัชนี ความคลาดเคลื่อนในการประมาณค่าพารามิเตอร์ (Root Mean Square Error of Approximation)

1.2 สัญลักษณ์ที่ใช้แทนตัวแปรแฝง

MANAGE	หมายถึง การบริหารจัดการแบบมีส่วนร่วม
SUCCESS	หมายถึง ปัจจัยแห่งความสำเร็จ
DHAMMA	หมายถึง ธรรมะเพื่อการปฏิบัติ
EFFECTIVE	หมายถึง ประสิทธิภาพหรือวิธีการบริหารจัดการศูนย์ฯ

1.3 สัญลักษณ์ที่ใช้แทนตัวแปรสังเกตได้

MANAGEA	หมายถึง ร่วมคิด
MANAGEB	หมายถึง ร่วมตัดสินใจ
MANAGEC	หมายถึง ร่วมวางแผน
MANAGED	หมายถึง ร่วมดำเนินการ
MANAGEE	หมายถึง ร่วมติดตามประเมินผล
MANAGEF	หมายถึง ร่วมรับผลประโยชน์
SUCCESSA	หมายถึง ความรู้
SUCCESSB	หมายถึง ทักษะ/ความสามารถ
SUCCESSC	หมายถึง เทคโนโลยี/นวัตกรรม
SUCCESSD	หมายถึง เกษตรกรต้นแบบ
SUCCESE	หมายถึง ทุน
SUCCESSF	หมายถึง การตลาด

DHAMA	หมายถึง มีใจรัก (ฉันทะ)
DHAMB	หมายถึง พากเพียรทำ (วิริยะ)
DHAMC	หมายถึง นำจิตฝึกฝน (จิตตะ)
DHAMD	หมายถึง ใช้ปัญญาทบทวน (วิมังสา)
EFFECA	หมายถึง พัฒนาไม่หยุดยั้ง
EFFECB	หมายถึง รวมพลังสร้างเครือข่าย
EFFECC	หมายถึง ดำรงอยู่ไม่เสื่อมคลาย
EFFECD	หมายถึง ยอมรับได้ในสังคม

1.4 การนำเสนอผลการวิเคราะห์โมเดลสมการโครงสร้าง

ในการวิจัยครั้งนี้ ผู้วิจัยได้นำเสนอผลการวิเคราะห์ข้อมูลเชิงปริมาณที่ได้จากการใช้โปรแกรมสถิติขั้นสูง ซึ่งจะใช้ตัวเลขอารบิกประกอบการนำเสนอผลการวิเคราะห์ในตารางและแผนภาพ เพื่อความชัดเจนและสอดคล้องกับข้อมูลที่นำมาจากแผนภาพและ Print Out ของผลการวิเคราะห์ข้อมูล

2. ผลการวิเคราะห์ค่าสถิติเบื้องต้นของตัวแปรสังเกตได้ในโมเดลประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนเกษตร

ผลการวิเคราะห์ค่าสถิติพื้นฐานของตัวแปรที่ใช้ในแต่ละโมเดล มีตัวบ่งชี้ทั้งหมด 20 ตัว ที่ใช้วัดตัวแปรแฝง 4 ตัวแปร คือ การบริหารจัดการแบบมีส่วนร่วม (MANAGE) ปัจจัยแห่งความสำเร็จ (SUCCESS) ประสิทธิผลพุทธวิธีการบริหารจัดการศูนย์ (EFFECTIVE) ตัวแปรส่งผ่านคือ ธรรมะเพื่อการปฏิบัติ (DHAMMA) มีจุดมุ่งหมายเพื่อศึกษาลักษณะการกระจายและการแจกแจงของตัวแปรสังเกตได้แต่ละตัว ค่าสถิติที่ใช้ได้แก่ ค่าเฉลี่ย ($\bar{X}$) ส่วนเบี่ยงเบนมาตรฐาน (S.D.) คะแนนต่ำสุด (Min) คะแนนสูงสุด (Max) สัมประสิทธิ์การกระจาย (C.V.) ค่าความเบ้ (Sk) และค่าความโด่ง (Ku) โดยแยกวิเคราะห์ผล แต่ละตัวแปรดังต่อไปนี้

เมื่อพิจารณาตัวแปรการบริหารจัดการแบบมีส่วนร่วม (MANAGE) พบว่า โดยภาพรวมตัวแปรการบริหารจัดการแบบมีส่วนร่วม อยู่ในระดับปานกลาง ($\bar{X} = 3.02$) ซึ่งเมื่อพิจารณาเป็นรายด้านพบว่า การร่วมรับผลประโยชน์ มีค่าเฉลี่ยสูงที่สุด ($\bar{X} = 3.12$) รองลงมาคือ การร่วมคิด ($\bar{X} = 3.04$) การร่วมวางแผน ($\bar{X} = 3.02$) การร่วมติดตามประเมินผล ($\bar{X} = 2.98$) การร่วมดำเนินการ ($\bar{X} = 2.97$) และการร่วมตัดสินใจ ($\bar{X} = 2.95$) ตามลำดับ เมื่อพิจารณาค่าสัมประสิทธิ์การกระจาย (C.V.) ของตัวแปรพบว่า ตัวแปรมีการกระจายไม่ต่างกันมาก โดยมีค่าอยู่ระหว่างร้อยละ 34.24 – 39.07 เมื่อพิจารณาค่าความเบ้

(Sk) ของตัวแปร พบว่า ตัวแปรร่วมคิด ร่วมตัดสินใจ และร่วมติดตามประเมินผลมีการแจกแจงในลักษณะเบ้ขวา (ค่าความเบ้เป็นบวก) แสดงว่าข้อมูลตัวแปรต่ำกว่าค่าเฉลี่ยและตัวแปรร่วมวางแผน ร่วมดำเนินการ และร่วมรับผลประโยชน์มีการแจกแจงในลักษณะเบ้ซ้าย (ค่าความเบ้เป็นลบ) แสดงว่า ข้อมูลของตัวแปรทุกตัวสูงกว่าค่าเฉลี่ย เมื่อพิจารณาค่าความโด่ง (Ku) พบว่า ตัวแปรทุกตัวมีการแจกแจงของข้อมูลในลักษณะเตี้ยแบนกว่าโค้งปกติ (ค่าความโด่งน้อยกว่า 0) แสดงว่ามีการกระจายของข้อมูลมาก

เมื่อพิจารณาตัวแปรปัจจัยแห่งความสำเร็จ (SUCCESS) พบว่า โดยภาพรวมปัจจัยแห่งความสำเร็จอยู่ในระดับปานกลาง ($\bar{X} = 3.46$) ซึ่งเมื่อพิจารณาเป็นรายด้าน พบว่า ด้านเกษตรกรต้นแบบมีค่าเฉลี่ยสูงสุด ($\bar{X} = 3.62$) รองลงมาคือ ด้านความรู้ ($\bar{X} = 3.54$) และด้านทักษะ/ความสามารถ ($\bar{X} = 3.50$) ด้านเทคโนโลยี/นวัตกรรม ($\bar{X} = 3.43$) ด้านการตลาด ($\bar{X} = 3.31$) และด้านทุน ($\bar{X} = 3.26$) ตามลำดับ เมื่อพิจารณาค่าสัมประสิทธิ์การกระจาย (C.V.) ของตัวแปรพบว่า ตัวแปรมีการกระจายไม่ต่างกันมาก โดยมีค่าอยู่ระหว่างร้อยละ 27.90 – 31.72 เมื่อพิจารณาค่าความเบ้ (Sk) ของตัวแปร พบว่า ตัวแปรทุกตัวมีการแจกแจงในลักษณะเบ้ซ้าย (ค่าความเบ้เป็นลบ) แสดงว่า ข้อมูลของตัวแปรทุกตัวสูงกว่าค่าเฉลี่ย เมื่อพิจารณาค่าความโด่ง (Ku) พบว่า ตัวแปรทุกตัวมีโค้งการแจกแจงของข้อมูลในลักษณะเตี้ยแบนกว่าโค้งปกติ (ค่าความโด่งน้อยกว่า 0) แสดงว่ามีการกระจายของข้อมูลมาก

เมื่อพิจารณาตัวแปรด้านธรรมะเพื่อการปฏิบัติ (DHAMMA) พบว่า โดยภาพรวมตัวแปรธรรมะเพื่อการปฏิบัติอยู่ในระดับค่อนข้างสูง ($\bar{X} = 3.59$) ซึ่งเมื่อพิจารณาเป็นรายด้าน พบว่า ด้านมีใจรัก (ฉันทะ) สูงที่สุด ($\bar{X} = 3.68$) รองลงมาคือ พากเพียรทำ (วิริยะ) ($\bar{X} = 3.60$) นำจิตฝึกไฟ (จิตตะ) ($\bar{X} = 3.60$) และ ใช้ปัญญาทบทวน (วิมังสา) ($\bar{X} = 3.44$) ตามลำดับ เมื่อพิจารณาค่าสัมประสิทธิ์การกระจาย (C.V.) ของตัวแปร พบว่า ตัวแปรมีการกระจายไม่ต่างกันมาก โดยมีค่าอยู่ระหว่างร้อยละ 29.25 - 30.52 เมื่อพิจารณาค่าความเบ้ (Sk) ของตัวแปร พบว่า ตัวแปรทุกตัวมีการแจกแจงในลักษณะเบ้ซ้าย (ค่าความเบ้เป็นลบ) แสดงว่า ข้อมูลของตัวแปรทุกตัวสูงกว่าค่าเฉลี่ย เมื่อพิจารณาค่าความโด่ง (Ku) พบว่า ตัวแปรทุกตัวมีโค้งการแจกแจงของข้อมูลในลักษณะเตี้ยแบนกว่าโค้งปกติ (ค่าความโด่งน้อยกว่า 0) แสดงว่ามีการกระจายของข้อมูลมาก

เมื่อพิจารณาตัวแปรด้านประสิทธิผลพฤติกรรมการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร (EFFECTIVE) พบว่า โดยภาพรวมตัวแปรประสิทธิผลพฤติกรรมการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตรอยู่ในระดับปานกลาง ($\bar{X} = 3.48$) ซึ่งเมื่อพิจารณาเป็นรายด้าน พบว่า ด้านพัฒนาไม่หยุดยั้งสูงที่สุด ($\bar{X} = 3.51$) รองลงมาคือ ดำรงอยู่ไม่เสื่อมคลาย ($\bar{X} =$

3.49) ยอมรับได้ในสังคม ($\bar{X} = 3.48$) และรวมพลังสร้างเครือข่าย ($\bar{X} = 3.45$) ตามลำดับ เมื่อพิจารณาค่าสัมประสิทธิ์การกระจาย (C.V.) ของตัวแปร พบว่า ตัวแปรมีการกระจายไม่ต่างกันมาก โดยมีค่าอยู่ระหว่างร้อยละ 30.37 – 31.61 เมื่อพิจารณาค่าความเบ้ (Sk) ของตัวแปร พบว่า ตัวแปรทุกตัว มีการแจกแจงในลักษณะเบ้ซ้าย (ค่าความเบ้เป็นลบ) แสดงว่ามีข้อมูลของตัวแปรสูงกว่าค่าเฉลี่ย เมื่อพิจารณาค่าความโด่ง (Ku) พบว่า ตัวแปรทุกตัวมีค่าโด่งการแจกแจงของข้อมูลในลักษณะเตี้ยแบนกว่าโค้งปกติ (ค่าความโด่งน้อยกว่า 0) แสดงว่ามีการกระจายของข้อมูลมาก แสดงรายละเอียดของผลการวิเคราะห์ข้อมูลได้ดังตารางที่ 10.19

ตารางที่ 10.19 ค่าสถิติเบื้องต้นของตัวแปรสังเกตได้ในโมเดลประสิทธิภาพพฤติกรรมการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร

(N=520)

ตัวแปร	$\bar{X}$	ระดับ	S.D.	Min	Max	C.V.	Sk	Ku
MANAGE	3.04	ปานกลาง	1.07	1.12	5.00	35.19	0.07	-1.14
MANAGEA	2.95	ปานกลาง	1.01	1.38	5.00	34.24	0.33	-1.17
MANAGEB	3.01	ปานกลาง	1.14	1.00	5.00	37.87	0.09	-1.09
MANAGEC	2.97	ปานกลาง	1.15	1.00	5.00	38.72	-0.04	-1.07
MANAGED	2.98	ปานกลาง	1.13	1.00	5.00	37.92	-0.03	-1.07
MANAGEE	3.12	ปานกลาง	1.19	1.00	5.00	38.14	0.06	-1.08
MANAGEF	3.02	ปานกลาง	1.18	1.00	5.00	39.07	-0.18	-1.14
SUCCESS	3.46	ปานกลาง	0.96	1.00	5.00	27.75	-0.52	-0.45
SUCCESSA	3.54	ค่อนข้างสูง	1.08	1.00	5.00	30.50	-0.43	-0.85
SUCCESSB	3.50	ค่อนข้างสูง	0.99	1.00	5.00	28.28	-0.62	-0.23
SUCCESSC	3.43	ปานกลาง	1.03	1.00	5.00	30.02	-0.52	-0.39
SUCCESSD	3.62	ค่อนข้างสูง	1.01	1.00	5.00	27.90	-0.57	-0.42
SUCSESSE	3.26	ปานกลาง	1.03	1.00	5.00	31.59	-0.12	-1.03
SUCCESSF	3.31	ปานกลาง	1.05	1.00	5.00	31.72	-0.30	-0.94
DHAMMA	3.59	ค่อนข้างสูง	1.05	1.00	5.00	29.25	-0.75	-0.45
DHAMA	3.68	ค่อนข้างสูง	1.12	1.00	5.00	30.43	-0.74	-0.49
DHAMB	3.60	ค่อนข้างสูง	1.08	1.00	5.00	30.00	-0.63	-0.48
DHAMC	3.60	ค่อนข้างสูง	1.07	1.00	5.00	29.72	-0.67	-0.51
DHAMD	3.44	ปานกลาง	1.05	1.00	5.00	30.52	-0.57	-0.79

ตารางที่ 10.19 ค่าสถิติเบื้องต้นของตัวแปรสังเกตได้ในโมเดลประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร (ต่อ)

(N=520)

ตัวแปร	$\bar{X}$	ระดับ	S.D.	Min	Max	C.V.	Sk	Ku
EFFECTIVE	3.48	ปานกลาง	1.06	1.00	5.00	30.46	-0.47	-0.79
EFFECA	3.51	ค่อนข้างสูง	1.09	1.00	5.00	31.05	-0.49	-0.80
EFFECB	3.45	ปานกลาง	1.07	1.00	5.00	31.01	-0.47	-0.77
EFFECC	3.49	ปานกลาง	1.06	1.00	5.00	30.37	-0.42	-0.70
EFFECD	3.48	ปานกลาง	1.10	1.00	5.00	31.61	-0.35	-0.83

3. ผลการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรสังเกตได้เพื่อใช้สร้างเมทริกซ์สหสัมพันธ์ในการวิเคราะห์โมเดลประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร

ความสัมพันธ์ระหว่างตัวแปรการบริหารจัดการแบบมีส่วนร่วม (MANAGE) ปัจจัยแห่งความสำเร็จ (SUCCESS) ธรรมะเพื่อการปฏิบัติ (DHAMMA) และ ประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร (EFFECTIVE) มีรายละเอียดผลการวิเคราะห์ ดังนี้

โมเดลที่แสดงอิทธิพลของการบริหารจัดการแบบมีส่วนร่วม ปัจจัยแห่งความสำเร็จที่มีต่อประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร โดยมีธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน

ผลการวิเคราะห์ พบว่า ค่าสถิติ Bartlett's Test of Sphericity ซึ่งเป็นสถิติทดสอบสมมติฐานว่า เมทริกซ์สหสัมพันธ์เป็นเมทริกซ์เอกลักษณ์ (Identity matrix) หรือไม่มีค่าสถิติทดสอบเท่ากับ 20227.041 ($p=.000$) แสดงว่า เมทริกซ์สหสัมพันธ์ระหว่างตัวแปรสังเกตได้ทั้งหมดของกลุ่มตัวอย่างแตกต่างจากเมทริกซ์เอกลักษณ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 และค่าดัชนี ไกเซอร์ - ไมเยอร์ - ออลคิน (Kaiser-Meyer-Olkin Measure of Sampling Adequacy: KMO) มีค่าเท่ากับ .918 โดยมีค่าเข้าใกล้ 1 แสดงว่าตัวแปรในข้อมูลชุดนี้มีความสัมพันธ์กันเหมาะสมที่จะนำไปวิเคราะห์โมเดลลิสรต่อไป

เมื่อพิจารณาความสัมพันธ์ของตัวแปรสังเกตได้จำนวน 20 ตัวแปร พบว่าความสัมพันธ์ระหว่างตัวแปรที่มีค่าแตกต่างจากศูนย์อย่างมีนัยสำคัญทางสถิติ ($p<.01$) มีจำนวน 190 คู่ มีค่าพิสัยสัมประสิทธิ์สหสัมพันธ์อยู่ในช่วง .411 - .960

เมื่อพิจารณาความสัมพันธ์ของตัวแปรสังเกตได้ พบว่า ทุกตัวมีความสัมพันธ์อย่างมีนัยสำคัญทางสถิติ ($p < .01$) และเป็นความสัมพันธ์ทางบวก แสดงว่า ความสัมพันธ์ของตัวแปรเป็นไปในทิศทางเดียวกัน โดยตัวแปรที่มีความสัมพันธ์ทางบวกสูงสุด คือ พัฒนาไม่หยุดยั้ง (EFFECA) และรวมพลังสร้างเครือข่าย (EFFECB) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .960 แสดงว่า เมื่อประสิทธิผลพฤติกรรมการบริหารจัดการศูนย์ฯด้านพัฒนาไม่หยุดยั้งเพิ่มขึ้น ประสิทธิผลพฤติกรรมการจัดการศูนย์ฯ ด้านรวมพลังสร้างเครือข่ายก็เพิ่มขึ้นด้วย และตัวแปรที่มีความสัมพันธ์ทางบวกรองลงมาคือ ประสิทธิผลพฤติกรรมการจัดการศูนย์ฯด้านดำรงอยู่ไม่เสื่อมคลาย (EFFECC) และ ประสิทธิผลพฤติกรรมการจัดการศูนย์ฯด้านยอมรับได้ในสังคม (EFFECD) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .959 แสดงว่า เมื่อประสิทธิผลพฤติกรรมการจัดการศูนย์ฯด้านดำรงอยู่ไม่เสื่อมคลายสูงขึ้น ประสิทธิผลพฤติกรรมการจัดการศูนย์ฯด้านยอมรับได้ในสังคมก็สูงขึ้นด้วย

เมื่อพิจารณาความสัมพันธ์ของตัวแปรสังเกตได้ระหว่างกลุ่มตัวแปรด้านเดียวกันมีรายละเอียดดังต่อไปนี้

ตัวแปรสัมพันธภาพของการบริหารจัดการแบบมีส่วนร่วม (MANAGE) พบว่า มีค่าพิสัยสัมประสิทธิ์สหสัมพันธ์อยู่ในช่วง .825 ถึง .948 โดยตัวแปรที่มีความสัมพันธ์กันสูงสุดคือ ร่วมคิด (MANAGEA) และร่วมตัดสินใจ (MANAGEB) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .825 แสดงว่า เมื่อการบริหารจัดการแบบร่วมคิดสูงขึ้น การบริหารจัดการแบบร่วมตัดสินใจก็สูงขึ้นด้วย ส่วนตัวแปรที่มีความสัมพันธ์กันต่ำสุดคือ ร่วมติดตามประเมินผล (MANAGEE) และ ร่วมรับผลประโยชน์ (MANAGEF) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .825

ตัวแปรปัจจัยแห่งความสำเร็จ (SUCCESS) พบว่า มีค่าพิสัยสัมประสิทธิ์สหสัมพันธ์อยู่ในช่วง .747 ถึง .939 โดยตัวแปรที่มีความสัมพันธ์กันสูงสุดคือ ทักษะ/ความสามารถ (SUCCESB) และ เทคโนโลยี/นวัตกรรม (SUCCESSC) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .939 เมื่อทักษะ/ความสามารถสูงขึ้น เทคโนโลยี/นวัตกรรม ก็สูงขึ้นด้วย ส่วนตัวแปรที่มีความสัมพันธ์กันต่ำสุดคือ ความรู้ (SUCCESSA) และการตลาด (SUCCESSF) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .747

ตัวแปรธรรมะเพื่อการปฏิบัติ (DHAMMA) พบว่า มีค่าพิสัยสัมประสิทธิ์สหสัมพันธ์อยู่ในช่วง .860 ถึง .947 โดยตัวแปรที่มีความสัมพันธ์กันสูงสุดคือ พากเพียรทำ

(วิริยะ) (DHAMB) และ นำจิตฝึกไฟ (จิตตะ) (DHAMC) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .947 แสดงว่า เมื่อปากเพียรทำ (วิริยะ) สูงขึ้น นำจิตฝึกไฟ (จิตตะ) ก็จะสูงขึ้นด้วย ส่วนตัวแปรที่มีความสัมพันธ์กันต่ำที่สุดคือ คือ ปากเพียรทำ (วิริยะ) (DHAMB) และใช้ปัญญาทบทวน (วิมังสา) (DHAMD) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .860

ตัวแปรประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์ฯ (EFFECTIVE) พบว่า มีค่าพิสัยสัมประสิทธิ์สหสัมพันธ์อยู่ในช่วง .911 ถึง .960 โดยมีตัวแปรที่มีความสัมพันธ์กันสูงสุดคือ พัฒนาไม่หยุดยั้ง (EFFECA) และ รวมพลังสร้างเครือข่าย (EFFECB) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .960 แสดงว่า เมื่อพัฒนาไม่หยุดยั้งสูงขึ้น รวมพลังสร้างเครือข่าย ก็จะสูงขึ้นด้วย ส่วนตัวแปรที่มีความสัมพันธ์กันต่ำที่สุดคือ พัฒนาไม่หยุดยั้ง (EFFECA) และ ยอมรับได้ในสังคม (EFFECD) โดยมีขนาดความสัมพันธ์อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 เท่ากับ .911 สามารถแสดงผลการวิเคราะห์ความสัมพันธ์ของตัวแปรสังเกตได้ในโมเดลประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ได้ดังตารางที่ 10.20

ตารางที่ 10.20 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และค่าสัมประสิทธิ์สหสัมพันธ์แบบเพียร์สันของตัวแปรสังเกตได้ในโมเดล ประสิทธิภาพ
พฤติกรรมวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ที่มีตัวแปรอิสระเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน

ตัวแปร	MANAGEA	MANAGEB	MANAGEC	MANAGED	MANAGEE	MANAGEF	SUCCESSA	SUCCESSB	SUCCESSC	SUCCESSD	SUCCESE	SUCCESSF
MANAGEA	1.000											
MANAGEB	.948**	1.000										
MANAGEC	.909**	.910**	1.000									
MANAGED	.932**	.945**	.918**	1.000								
MANAGEE	.902**	.896**	.914**	.935**	1.000							
MANAGEF	.835**	.858**	.860**	.844**	.825**	1.000						
SUCCESSA	.585**	.603**	.648**	.581**	.528**	.679**	1.000					
SUCCESSB	.575**	.591**	.643**	.597**	.535**	.678**	.932**	1.000				
SUCCESSC	.541**	.579**	.583**	.575**	.532**	.654**	.894**	.939**	1.000			
SUCCESSD	.578**	.610**	.621**	.602**	.552**	.653**	.851**	.873**	.874**	1.000		
SUCCESE	.554**	.608**	.668**	.586**	.572**	.637**	.796**	.832**	.843**	.806**	1.000	
SUCCESSF	.565**	.626**	.636**	.572**	.529**	.644**	.747**	.773**	.776**	.757**	.796**	1.000
DHAMA	.606**	.585**	.566**	.581**	.533**	.615**	.675**	.664**	.651**	.677**	.551**	.638**
DHAMB	.526**	.516**	.466**	.505**	.441**	.540**	.625**	.613**	.622**	.682**	.528**	.562**
DHAMC	.582**	.563**	.496**	.558**	.489**	.565**	.656**	.629**	.644**	.707**	.553**	.604**
DHAMD	.603**	.591**	.584**	.567**	.552**	.623**	.658**	.619**	.601**	.667**	.564**	.600**
EFFECA	.562**	.567**	.563**	.544**	.505**	.680**	.864**	.829**	.812**	.770**	.687**	.668**
EFFECB	.549**	.543**	.528**	.529**	.483**	.661**	.842**	.828**	.807**	.762**	.678**	.708**
EFFECC	.603**	.603**	.555**	.569**	.519**	.682**	.849**	.837**	.831**	.766**	.698**	.700**
EFFECD	.525**	.526**	.537**	.502**	.449**	.646**	.867**	.843**	.840**	.759**	.724**	.735**
MEAN	3.045	2.949	3.012	2.979	2.981	3.116	3.537	3.500	3.431	3.618	3.258	3.309
SD	1.009	1.137	1.146	1.131	1.199	1.184	1.084	0.994	1.028	1.013	1.025	1.055

ตารางที่ 10.20 ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และค่าสัมประสิทธิ์สหสัมพันธ์แบบเพียร์สันของตัวแปรสังเกตได้ในโมเดล ประสิทธิภาพพุทธ
วิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ที่มีตัวแปรธรรมชาติเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน (ต่อ)

ตัวแปร	DHAMA	DHAMB	DHAMC	DHAMD	EFFECA	EFFECB	EFFECC	EFFECD
DHAMA	1.000							
DHAMB	.925**	1.000						
DHAMC	.939**	.947**	1.000					
DHAMD	.902**	.806**	.913**	1.000				
EFFECA	.733**	.692**	.740**	.758**	1.000			
EFFECB	.748**	.702**	.676**	.754**	.960**	1.000		
EFFECC	.765**	.726**	.778**	.743**	.925**	.956**	1.000	
EFFECD	.724**	.683**	.730**	.709**	.911**	.939**	.959**	1.000
MEAN	3.684	3.601	3.601	3.441	3.509	3.451	3.495	3.486
SD	1.121	1.081	1.068	1.047	1.091	1.070	1.068	1.100
Bartlett's Test of Sphericity=20227.041 df=190 p=.000 Kaiser-Meyer - Olkin Measure of Sampling Adequacy=.918								

**p < .01

4. ผลการวิเคราะห์ความสอดคล้องกับข้อมูลเชิงประจักษ์ของโมเดล ประสิทธิผลพุทธิวิธีบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร

โมเดลนี้มีตัวแปรแฝง 4 ตัวแปร คือ การบริหารจัดการแบบมีส่วนร่วม (MANAGE) ปัจจัยแห่งความสำเร็จ (SUCCESS) ธรรมะเพื่อการปฏิบัติ (DHAMMA) และประสิทธิผลพุทธิวิธีบริหารจัดการศูนย์ฯ (EFFECTIVE) โดยตัวแปรสังเกตได้ที่ใช้ในการวิเคราะห์ข้อมูลทั้งหมด 20 ตัวแปร

การทดสอบความสอดคล้องของโมเดลประสิทธิผลพุทธิวิธีบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่านนี้ ในครั้งแรก พบว่า โมเดลไม่สอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ โดยพิจารณาจากค่าไค-สแควร์ มีค่าเท่ากับ 2154.11 ที่ค่าองศาอิสระเท่ากับ 164 และค่าความน่าจะเป็น (p) เท่ากับ .000 ดัชนีวัดความกลมกลืน (GFI) มีค่าเท่ากับ .710 ค่าดัชนีความกลมกลืนที่ปรับแก้แล้ว (AGFI) มีค่าเท่ากับ .620 ค่าดัชนีรากของกำลังสองเฉลี่ยของส่วนที่เหลือ (RMR) มีค่าเท่ากับ .061 ค่าดัชนีรากกำลังสองเฉลี่ยของค่าความแตกต่างโดยประมาณ (RMSEA) มีค่าเท่ากับ .153 และค่าเฉลี่ยเหลือในรูปคะแนนมาตรฐานระหว่างตัวแปรสูงสุด (Largest Standardized Residuals) เท่ากับ 12.36

จากผลการวิเคราะห์ดังกล่าว ผู้วิจัยจึงปรับโมเดลโดยยอมให้ความคลาดเคลื่อนสัมพันธ์กันได้ ซึ่งเป็นการผ่อนคลายข้อตกลงเบื้องต้นจากข้อตกลงเบื้องต้นในสถิติวิเคราะห์ดั้งเดิมที่กำหนดว่าทอมความคลาดเคลื่อนต้องไม่สัมพันธ์กัน เป็นข้อตกลงเบื้องต้นในสถิติวิเคราะห์ด้วย SEM ซึ่งกำหนดให้มีการนำทอมความคลาดเคลื่อนมาใช้ในการวิเคราะห์ข้อมูล และทอมความคลาดเคลื่อนมีความสัมพันธ์กันตามสภาพความเป็นจริงของปรากฏการณ์ธรรมชาติ ผลการปรับโมเดล จะได้ค่าขนาดอิทธิพลและค่าความสัมพันธ์ระหว่างตัวแปรในโมเดลที่ถูกต้องตรงกับความเป็นจริงมากขึ้น ผู้วิจัยพิจารณาปรับโมเดลจากดัชนีดัดแปลงโมเดล (Modification Indices) และได้ปรับโมเดลจำนวน 135 เส้นทาง โดยได้ปรับ 1) เส้นทาง Theta-Delta (TD) และ 2) เส้นทาง Theta-Epsilon (TE) และผลการปรับโมเดลทำให้ได้โมเดลประสิทธิผลพุทธิวิธีบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน ที่สอดคล้องกับข้อมูลเชิงประจักษ์

เมื่อพิจารณาผลการวิเคราะห์โมเดลประสิทธิผลพุทธิวิธีบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตร ที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน พบว่า โมเดลมีความสอดคล้องกับข้อมูลเชิงประจักษ์ พิจารณาจากค่าสถิติที่ใช้ตรวจสอบความสอดคล้องระหว่างโมเดลกับข้อมูลเชิงประจักษ์ ได้แก่ ค่าไค-สแควร์ มีค่าเท่ากับ 41.01 องศาอิสระเท่ากับ

29 ความน่าจะเป็น (p) เท่ากับ .069 นั่นคือ ค่าไค-สแควร์ แตกต่างจากศูนย์อย่างไม่มีนัยสำคัญ แสดงว่า ยอมรับสมมติฐานหลักที่ว่า โมเดลประสิทธิผลพฤติกรรมวิธีการบริหารจัดการศูนย์เรียนรู้ ชุมชนการเกษตร ที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน ที่พัฒนาขึ้นสอดคล้อง กลมกลืนกับข้อมูลเชิงประจักษ์ ซึ่งสอดคล้องกับผลการวิเคราะห์ค่าดัชนีวัดความกลมกลืน (GFI) มีค่าเท่ากับ .992 ค่าดัชนีวัดความกลมกลืนที่ปรับแก้แล้ว (AGFI) มีค่าเท่ากับ .943 ซึ่งมี ค่าเข้าใกล้ 1 และค่าดัชนีรากของกำลังสองเฉลี่ยของส่วนที่เหลือ (RMR) มีค่าเท่ากับ .013 ซึ่งเข้าใกล้ศูนย์ ค่าดัชนีรากกำลังสองเฉลี่ยของค่าความแตกต่างโดยประมาณ (RMSEA) มี ค่าเท่ากับ .028 ซึ่งเข้าใกล้ศูนย์ สนับสนุนว่าโมเดลการวิจัยมีความสอดคล้องกับข้อมูลเชิง ประจักษ์

เมื่อพิจารณาค่าความเที่ยงของตัวแปรสังเกตได้ พบว่า ตัวแปรสังเกตได้มีค่าความ เที่ยงอยู่ระหว่าง .666 ถึง 1.130 โดยตัวแปรที่มีค่าความเที่ยงสูงสุดคือ การบริหารจัดการแบบ มีส่วนร่วมด้านร่วมรับผลประโยชน์ (MANAGEF) มีค่าความเที่ยงเท่ากับ 1.130 รองลงมาคือ การบริหารแบบมีส่วนร่วมด้านร่วมคิด (MANAGEA) มีค่าความเที่ยงเท่ากับ 1.008 การ บริหารจัดการแบบมีส่วนร่วมด้าน ร่วมตัดสินใจ (MANAGEB) มีค่าความเที่ยงเท่ากับ 1.004 ปัจจัยแห่งความสำเร็จด้านความรู้ (SUCESSA) มีค่าความเที่ยงเท่ากับ 1.000 และตัวแปรที่มี ค่าความเที่ยงต่ำสุด คือ ปัจจัยแห่งความสำเร็จด้านทุน (MANAGEE) มีค่าความเที่ยงเท่ากับ .666 ในภาพรวมค่าความเที่ยงของตัวแปรสังเกตได้มีค่าอยู่ในระดับปานกลางถึงระดับสูงมาก

เมื่อพิจารณาค่าสัมประสิทธิ์พยากรณ์ (R-SQUARE) ของสมการโครงสร้างตัว แปรภายในแฝง พบว่า ธรรมะเพื่อการปฏิบัติ (DHAMMA) มีค่าสัมประสิทธิ์การพยากรณ์ เท่ากับ .502 แสดงว่า ตัวแปรภายในโมเดลประกอบด้วย การบริหารจัดการแบบมีส่วนร่วม (MANAGE) และปัจจัยแห่งความสำเร็จ (SUCCESS) สามารถอธิบายความแปรปรวนของ ธรรมะเพื่อการปฏิบัติ (DHAMMA) ได้ร้อยละ 50.20 ส่วนประสิทธิผลพฤติกรรมวิธีการบริหาร จัดการศูนย์ฯ (EFFEC) มีค่าสัมประสิทธิ์การพยากรณ์เท่ากับ .850 แสดงว่า ตัวแปรภายใน โมเดล คือ การบริหารจัดการแบบมีส่วนร่วม (MANAGE) ปัจจัยแห่งความสำเร็จ (SUCCESS) และธรรมะเพื่อการปฏิบัติ (DHAMMA) สามารถอธิบายความแปรปรวนของ ประสิทธิผลพฤติกรรมวิธีการบริหารจัดการศูนย์ฯ (EFFEC) ได้ร้อยละ 85.00

เมื่อพิจารณาเมตริกสหสัมพันธ์ระหว่างตัวแปรแฝง พบว่า ค่าพิสัยสัมประสิทธิ์ สหสัมพันธ์ระหว่างตัวแปรแฝงมีค่าอยู่ในช่วง .545 ถึง .878 โดยตัวแปรมีความสัมพันธ์แบบ ทิศทางเดียว (ค่าความสัมพันธ์เป็นบวก) ตัวแปรที่มีค่าสัมประสิทธิ์สหสัมพันธ์มากที่สุด คือ ปัจจัยแห่งความสำเร็จ (SUCCESS) และประสิทธิผลพฤติกรรมวิธีการบริหารจัดการศูนย์ฯ (EFFEC)

โดยมีค่าสัมประสิทธิ์สหสัมพันธ์เท่ากับ .878 แสดงว่า เมื่อปัจจัยแห่งความสำเร็จ (SUCCESS) มีมากขึ้น ประสิทธิภาพพฤติกรรมการจัดการศูนย์ฯ (EFFEC) ก็เพิ่มมากขึ้นด้วย และตัวแปรที่มีค่าสัมประสิทธิ์สหสัมพันธ์รองลงมาคือ ธรรมะเพื่อการปฏิบัติ (DHAMMA) และ ประสิทธิภาพพฤติกรรมการจัดการศูนย์ฯ (EFFEC) โดยมีค่าสัมประสิทธิ์สหสัมพันธ์เท่ากับ .795 ซึ่งเป็นความสัมพันธ์ที่อยู่ในระดับสูง

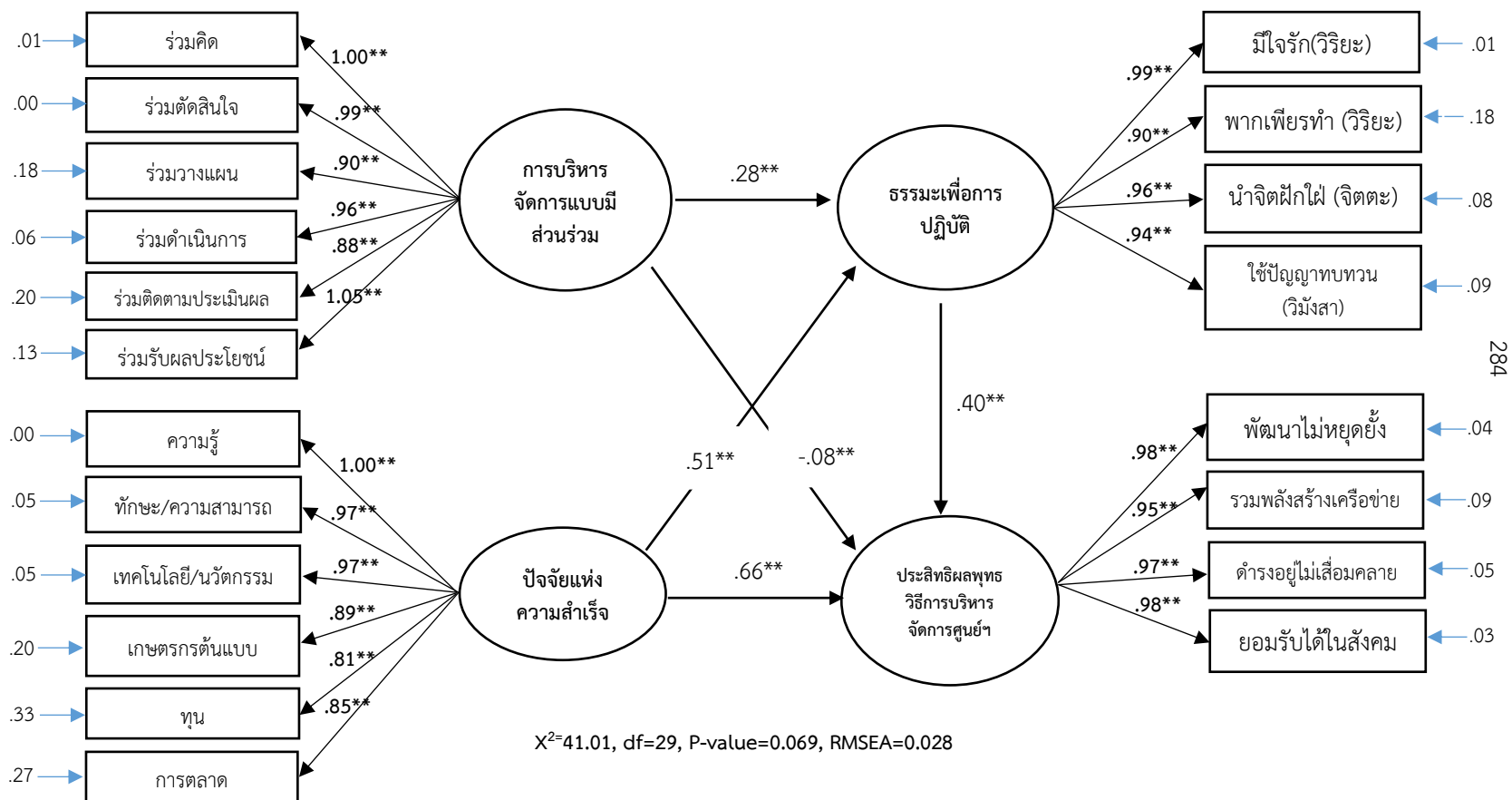
เมื่อพิจารณาอิทธิพลทางตรงและอิทธิพลทางอ้อมที่ส่งผลต่อตัวแปรประสิทธิผลพฤติกรรมการจัดการศูนย์ฯ (EFFEC) พบว่าตัวแปรดังกล่าว ได้รับอิทธิพลทางตรงจากการบริหารจัดการแบบมีส่วนร่วม (MANAGE) ปัจจัยแห่งความสำเร็จ (SUCCESS) และธรรมะเพื่อการปฏิบัติ (DHAMMA) โดยมีขนาดอิทธิพลเท่ากับ -.081, .658 และ .399 อย่างมีนัยสำคัญทางสถิติที่ระดับ .01 ตามลำดับ นอกจากนี้ ประสิทธิภาพพฤติกรรมการจัดการศูนย์ฯ (EFFEC) ยังได้รับอิทธิพลทางอ้อมจากการบริหารจัดการแบบมีส่วนร่วม (MANAGE) และ ปัจจัยแห่งความสำเร็จ (SUCCESS) ผ่านตัวแปรธรรมะเพื่อการปฏิบัติ (DHAMMA) โดยมีขนาดอิทธิพลเท่ากับ .109 และ .203 ตามลำดับอย่างมีนัยสำคัญทางสถิติที่ระดับ .01 ซึ่งจะเห็นว่าขนาดอิทธิพลทางตรงของปัจจัยแห่งความสำเร็จ (SUCCESS) มีขนาดสูงกว่าอิทธิพลทางอ้อม แสดงว่าธรรมะเพื่อการปฏิบัติ (DHAMMA) ไม่เป็นตัวแปรส่งผ่านที่จะทำให้ ประสิทธิภาพพฤติกรรมการจัดการศูนย์ฯ มากขึ้น แต่ในทางตรงกันข้ามจะเห็นว่าขนาดอิทธิพลทางอ้อมของการบริหารจัดการแบบมีส่วนร่วม (MANAGE) มีขนาดสูงกว่าอิทธิพลทางตรง แสดงว่า ธรรมะเพื่อการปฏิบัติ (DHAMMA) เป็นตัวแปรส่งผ่านที่ดีทำให้ประสิทธิภาพพฤติกรรมการจัดการศูนย์ฯ มากขึ้น นอกจากอิทธิพลทางตรงและทางอ้อมที่ส่งผลต่อ ประสิทธิภาพพฤติกรรมการจัดการศูนย์ฯ (EFFEC) แล้ว ยังมีตัวแปรธรรมะเพื่อการปฏิบัติ (DHAMMA) ที่ได้รับอิทธิพลทางตรงจาก การบริหารจัดการแบบมีส่วนร่วม (MANAGE) และ ปัจจัยแห่งความสำเร็จ (SUCCESS) โดยมีขนาดอิทธิพลเท่ากับ .273 และ .509 ตามลำดับ อย่างมีนัยสำคัญทางสถิติที่ระดับ .01

ตารางที่ 10.21 ค่าสถิติผลการวิเคราะห์แยกค่าสหสัมพันธ์ระหว่างตัวแปรแฝงและการวิเคราะห์อิทธิพลของโมเดลประสิทธิผลพุทธรูปวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตรที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน

ตัวแปรเหตุ	MANAGE			SUCCESS			DHAMMA		
ตัวแปรผล	TE	IE	DE	TE	IE	DE	TE	IE	DE
DHAMMA	.273** (.037)		.273** (.040)	.509** (.040)		.509** (.040)			
EFFECTIVE	.028 (.026)	.109** (.017)	-.081** (.025)	.861** (.039)	.203** (.024)	.658** (.036)	.399** (.035)		.399** (.036)
ค่าสถิติ	ไค-สแควร์=41.01 df=29 p=.069 GFI=.992 AGFI=.951 RMSEA=.028								
ตัวแปร	MANAGEA	MANAGEB	MANAGEC	MANAGED	MANAGEE	MANAGEF			
ความเที่ยง	1.008	1.004	.818	.938	.800	1.130			
ตัวแปร	SUCCESSA	SUCCESSB	SUCCESSC	SUCCESSD	SUCCESE	SUCCESSF			
ความเที่ยง	1.000	.949	.953	.803	0.666	0.730			
ตัวแปร	DHAMA		DHAMB		DHAMC		DHAMD		
ความเที่ยง	.994		.818		.922		.905		
ตัวแปร	EFFECA		EFFECB		EFFECC		EFFECD		
ความเที่ยง	.963		.909		.951		.969		
สมการโครงสร้างของตัวแปร				DHAMMA			EFFEC		
R SQUARE				.502			.850		
เมทริกซ์สหสัมพันธ์ระหว่างตัวแปรแฝง									
ตัวแปรแฝง	DHAMMA			EFFECTIVE			MANAGE		SUCCESS
DHAMMA	1.000								
EFFECTIVE	.795			1.000					
MANAGE	.580			.545			1.000		
SUCCESS	.674			.878			.601		1.000

หมายเหตุ: ตัวเลขในวงเล็บคือค่าความคลาดเคลื่อนมาตรฐาน, **p<.01

TE=ผลรวมอิทธิพล, IE=อิทธิพลทางอ้อม, DE=อิทธิพลทางตรง



แผนภาพที่ 10.5 โมเดลประสิทธิภาพผลพุทธวิธีการบริหารจัดการศูนย์เรียนรู้ชุมชนการเกษตรที่มีตัวแปรธรรมะเพื่อการปฏิบัติเป็นตัวแปรส่งผ่าน (ปรับโมเดลแล้ว)

สรุปได้ว่า SEM มีบทบาทสำคัญในงานวิจัยด้านรัฐประศาสนศาสตร์ เพราะช่วยให้นักวิจัยสามารถสร้างและยืนยันโมเดลที่ซับซ้อน ซึ่งไม่เพียงตรวจสอบโครงสร้างทางทฤษฎี แต่ยังสนับสนุนการกำหนดนโยบายและแนวทางบริหารงานภาครัฐได้อย่างมีหลักฐานเชิงประจักษ์ งานวิจัยของสุภัทรชัย สีสะใบ (2563) เป็นตัวอย่างที่ชัดเจนว่าการใช้ SEM สามารถสร้างแบบจำลองการบริหารที่มีความน่าเชื่อถือและสามารถนำไปปรับใช้ได้จริงในระดับปฏิบัติการ

สรุปท้ายบท

สถิติมีบทบาทสำคัญอย่างยิ่งต่อการพัฒนางานวิจัยด้านรัฐประศาสนศาสตร์ทั้งในมิติของการสร้างองค์ความรู้ การตรวจสอบสมมติฐานเชิงทฤษฎี และการนำผลลัพธ์ไปใช้ประโยชน์ในการบริหารจัดการและการกำหนดนโยบายสาธารณะ โดยสถิติทำหน้าที่เป็นเครื่องมือที่ช่วยให้นักวิจัยสามารถอธิบายปรากฏการณ์ทางสังคมและการเมืองได้อย่างมีระบบและตรวจสอบได้ตามหลักวิทยาศาสตร์

สถิติเชิงพรรณนาเป็นขั้นพื้นฐานที่ช่วยให้นักวิจัยสามารถอธิบายลักษณะทั่วไปของกลุ่มตัวอย่างและทัศนคติของประชาชนได้อย่างเป็นระบบ ผ่านการนำเสนอข้อมูลในรูปของค่าเฉลี่ย ร้อยละ ค่าเฉลี่ย และส่วนเบี่ยงเบนมาตรฐาน เครื่องมือนี้มีประโยชน์ในการทำให้อธิบายข้อมูลจำนวนมากที่ซับซ้อนถูกสรุปให้อยู่ในรูปแบบที่เข้าใจง่ายและสะท้อนภาพรวมได้ชัดเจน อันเป็นพื้นฐานในการตีความเชิงลึกต่อไป

ต่อมาคือสถิติเชิงอนุมาน ซึ่งมีบทบาทในการทดสอบสมมติฐาน อธิบายความสัมพันธ์ และทำนายปัจจัยที่ส่งผลต่อการบริหารงานภาครัฐ ไม่ว่าจะเป็นการใช้ t-test เพื่อเปรียบเทียบความแตกต่างระหว่างสองกลุ่ม การใช้ ANOVA เพื่อเปรียบเทียบความแตกต่างระหว่างหลายกลุ่ม การใช้ Correlation เพื่อตรวจสอบความสัมพันธ์ระหว่างตัวแปร หรือการใช้ Regression เพื่อทำนายและอธิบายอิทธิพลของตัวแปรอิสระต่อการเปลี่ยนแปลงของตัวแปรตาม การใช้สถิติเหล่านี้ช่วยให้ข้อค้นพบของงานวิจัยก้าวข้ามจากการเป็นเพียงการสังเกตไปสู่การสร้างข้อสรุปที่สามารถนำไปใช้เป็นหลักฐานเชิงประจักษ์ในการกำหนดนโยบายและการบริหารจัดการภาครัฐได้จริง

ในระดับที่ซับซ้อนยิ่งขึ้น งานวิจัยทางรัฐประศาสนศาสตร์สมัยใหม่มีการนำสถิติขั้นสูง เช่น การวิเคราะห์องค์ประกอบเชิงยืนยัน (CFA) และการวิเคราะห์สมการโครงสร้าง (SEM) มาใช้ในการตรวจสอบความตรงเชิงโครงสร้างของเครื่องมือวิจัยและสร้างโมเดลเชิงสาเหตุที่สะท้อนความสัมพันธ์ของตัวแปรทั้งที่สังเกตได้และตัวแปรแฝง สถิติขั้นสูงเหล่านี้ช่วยให้นักวิจัยสามารถอธิบายกลไกที่ซ่อนอยู่ภายใต้ปรากฏการณ์ทางสังคมและการบริหารได้อย่างลึกซึ้ง รวมทั้งทำให้สามารถสร้างโมเดลเชิงทฤษฎีที่สอดคล้องกับข้อมูลเชิงประจักษ์

ผลลัพธ์จากการใช้สถิติขั้นสูงจึงไม่เพียงยืนยันคุณภาพของเครื่องมือวัด แต่ยังเพิ่มความน่าเชื่อถือและความสมบูรณ์ให้กับงานวิจัย ตลอดจนสนับสนุนการกำหนดนโยบายที่มีฐานข้อมูลรองรับ

กล่าวโดยสรุป สถิติไม่ได้เป็นเพียงเครื่องมือสำหรับการวิเคราะห์ข้อมูลเชิงตัวเลขเท่านั้น แต่ยังเป็นสะพานเชื่อมระหว่างทฤษฎีและการปฏิบัติในทางรัฐประศาสนศาสตร์ ผลงานวิจัยจริงที่นำเสนอในบทนี้ยืนยันได้ว่าการเลือกใช้สถิติอย่างเหมาะสมในแต่ละระดับ ตั้งแต่สถิติเชิงพรรณนา อนุมาน ไปจนถึงขั้นสูง สามารถเพิ่มความน่าเชื่อถือของข้อค้นพบ และทำให้งานวิจัยด้านรัฐประศาสนศาสตร์มีคุณค่าเชิงปฏิบัติ สามารถนำไปใช้ในการแก้ไขปัญหาและพัฒนาสังคมได้อย่างยั่งยืน

บรรณานุกรม

1. ภาษาไทย

(1) หนังสือ:

- กฤติยา วงศ์ก้อม. (2545). *ระเบียบวิธีวิจัยทางสังคมศาสตร์*. กรุงเทพฯ : สถาบันราชภัฏสวนสุนันทา.
- ณัฐรฐนนท์ กานต์วิกุลธนา. (2565). *สถิติ .. การวิจัยทางสังคมศาสตร์* [หนังสืออิเล็กทรอนิกส์, โครงการ “กระบวนการจัดการศึกษาระดับบัณฑิตศึกษาเพื่อพัฒนานักวิจัยมืออาชีพ”]. กรุงเทพฯ: มหาวิทยาลัยสวนดุสิต.
- ทวีรัตน์ ศิวดุญ. (2539). *สถิติและความน่าจะเป็น*. กรุงเทพฯ: แมคกรอ-ฮิล.
- ทวีศักดิ์ นพเกษร. (2548). *วิธีการวิจัยเชิงคุณภาพ เล่ม 1 : คู่มือปฏิบัติการวิจัยประยุกต์เพื่อพัฒนาคน องค์กร ชุมชน สังคม*. นครราชสีมา : ชมรมพยาบาลชุมชนแห่งประเทศไทย.
- นงพรรณ พิริยานุพงศ์. (2546). *คู่มือวิจัยและพัฒนา*. นนทบุรี : โครงการสวัสดิการวิชาการ สถาบันพระบรมราชชนก.
- นงลักษณ์ วิรัชชัย. (2543). *พรมแดนความรู้ด้านการวิจัยและสถิติ*. ชลบุรี : วิทยาลัยการบริหารรัฐกิจ มหาวิทยาลัยบูรพา.
- นิภา ศรีโพธิ์โรจน์. (2531). *หลักการวิจัยเบื้องต้น*. พิมพ์ครั้งที่ 2. กรุงเทพฯ : บริษัท ศึกษาพร จำกัด.
- บุญเชิด ภิญโญนนตพงษ์. (2545). *การประเมินการเรียนรู้ที่เน้นผู้เรียนเป็นสำคัญ แนวคิดและวิธีการ*. พิมพ์ครั้งที่ 2. กรุงเทพฯ: ไทยวัฒนาพานิช.
- บุญธรรม กิจปรีดาบริสุทธิ์. (2534). *การวิจัย การวัดและประเมินผล*. กรุงเทพฯ : ภาควิชาศึกษาศาสตร์ คณะสังคมศาสตร์และมนุษยศาสตร์ มหาวิทยาลัยมหิดล.
- บุปผา ศิริรัศมี และคณะ. (2544). *จริยธรรมสำหรับการศึกษาวิจัยในคน*. เอกสารวิชาการ หมายเลข 258. นครปฐม: สถาบันวิจัยประชากรและสังคม มหาวิทยาลัยมหิดล.
- ประยูร อิมิวัตร์. (2566). *การวิจัยทางรัฐประศาสนศาสตร์แบบบูรณาการ*. เชียงราย: มหาวิทยาลัยราชภัฏเชียงราย.
- ปาริชาติ สถาปิตานนท์. (2546). *ระเบียบวิธีวิจัยการสื่อสาร*. พิมพ์ครั้งที่ 2. สำนักพิมพ์แห่งจุฬาลงกรณ์มหาวิทยาลัย.
- ผ่องพรรณ ตรัยมงคลกุล. (2543). *การวิจัยในชั้นเรียน*. กรุงเทพฯ : มหาวิทยาลัยเกษตรศาสตร์.

- พิชิต ฤทธิจรุญ. (2544). *ระเบียบวิธีการวิจัยทางสังคมศาสตร์*. กรุงเทพฯ : คณะครุศาสตร์ สถาบันราชภัฏพระนคร.
- พิชิต ทัศนาวณิช. (2543). *สถิติเชิงพรรณนาและการประยุกต์ใช้*. กรุงเทพฯ: สำนักพิมพ์ มหาวิทยาลัยธรรมศาสตร์.
- มนตรี สังข์ทอง. (2557). *หลักสถิติ*. กรุงเทพฯ: ซีเอ็ดดูเคชั่น.
- ล้วน สายยศ และอังคณา สายยศ. (2538). *การวิจัยทางการศึกษาแนวทางปฏิบัติ*. กรุงเทพฯ: สุวีริยาสาส์น.
- ศิริชัย กาญจนวาสี, ทวีวัฒน์ ปิตยานนท์, ดิเรก ศรีสุขโข. (2537). *สถิติสำหรับการวิจัยทางการศึกษา และพฤติกรรมศาสตร์*. กรุงเทพฯ: โรงพิมพ์มหาวิทยาลัยศรีนครินทรวิโรฒ.
- สมจิต วัฒนาชยากุล. (2546). *สถิติพื้นฐานสำหรับนักวิทยาศาสตร์*. กรุงเทพฯ: ประกายพริก.
- สมชาย วรกิจเกษมสกุล. (2554). *ระเบียบวิธีการวิจัยทางพฤติกรรมศาสตร์และสังคมศาสตร์*. อุดรธานี: อักษรศิลป์การพิมพ์.
- สิน พันธุ์พินิจ. (2547). *เทคนิคการวิจัยทางสังคมศาสตร์*. กรุงเทพฯ : วิทย์พัฒนา.
- สินธะวา ความดิษฐ์. (2522). *วิธีวิจัยด้านการท่องเที่ยว*. กรุงเทพฯ. โรงพิมพ์มหาวิทยาลัยธุรกิจบัณฑิต.
- สุชาติ ประสิทธิ์รัฐสินธุ์. (2544). *ระเบียบวิธีวิจัยทางสังคมศาสตร์*. กรุงเทพฯ: เพ็ญฟ้าพรินต์.
- สุวิมล ตรีภานันท์. (2546). *การวิจัยทางการศึกษา*. กรุงเทพฯ: สำนักพิมพ์แห่งจุฬาลงกรณ์มหาวิทยาลัย.

(2) คุชฌินิพนธ์ วิทยานิพนธ์:

- ชญานุช สามัญ. (2561). *การพัฒนาทรัพยากรมนุษย์ของเทศบาลเมืองลำตาเสา อำเภอวังน้อย จังหวัดพระนครศรีอยุธยา. (สารนิพนธ์ปริญญารัฐประศาสนศาสตรมหาบัณฑิต). พระนครศรีอยุธยา: มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย.*
- นัตยา อินทร์กรุงเก่า. (2566). *ปัจจัยที่ส่งผลต่อประสิทธิภาพในการปฏิบัติงานตามหลักพุทธธรรมของบุคลากรมหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย. (สารนิพนธ์ปริญญารัฐประศาสนศาสตรมหาบัณฑิต). พระนครศรีอยุธยา: มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย.*
- พงศันคร โภชากรณ์. (2566). *การบูรณาการหลักพุทธธรรมเพื่อการบริหารจัดการโครงการลงทะเบียนเพื่อสวัสดิการแห่งรัฐ. (คุชฌินิพนธ์ปริญญาปรัชญาดุษฎีบัณฑิต สาขาวิชารัฐประศาสนศาสตร์). พระนครศรีอยุธยา: มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย.*

พระมหาสมักร อติภโท. (2568). *การบูรณาการหลักพุทธธรรมเพื่อส่งเสริมความนิยมทาง
การเมืองของประชาชนที่มีต่อพรรคการเมืองในจังหวัดพระนครศรีอยุธยา.
(ดุष्ฎิณีพนธ์ปริญญาปรัชญาดุษฎีบัณฑิต สาขาวิชารัฐศาสตร์).*
พระนครศรีอยุธยา: มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย.

สุภัทรชัย สีสะใบ. (2563). *การพัฒนาโมเดลประสิทธิภาพพฤติกรรมการบริหารจัดการศูนย์
เรียนรู้ชุมชนการเกษตร. (ดุष्ฎิณีพนธ์ปริญญาปรัชญาดุษฎีบัณฑิต สาขาวิชารัฐ
ประศาสนศาสตร์).* พระนครศรีอยุธยา: มหาวิทยาลัยมหาจุฬาลงกรณราช
วิทยาลัย.

(3) บทความ:

เบญจา ยอดดำเนิน-แอ๊ดติกิจ. (2553). ประเด็นจริยธรรมการวิจัยในคนทางด้าน
สังคมศาสตร์. *วารสารวิทยาศาสตร์เกษตรศาสตร์*, 31(2), 290-301.

ปริญญา สิริอิตตะกุล. (2555). การวิเคราะห์ความแปรปรวนแบบทางเดียว: การวิจัยทาง
สังคมศาสตร์. *วารสารสหวิทยาการวิจัย: ฉบับบัณฑิตศึกษา*, 1(1), 13-23.

สมถวิล วิจิตรวรรณ. (2565). สถิติความสัมพันธ์: เลือกใช้อย่างไร. *วารสารมนุษยศาสตร์
และสังคมศาสตร์ มหาวิทยาลัยราชพฤกษ์*, 8(2), 1-15.

อัศวิน แสงพิกุล. (2556). จริยธรรมการวิจัย. *วารสารสุทธิปริทัศน์*, 27(81), 136-147.

(4) สื่ออิเล็กทรอนิกส์:

กรมวิชาการเกษตร. (2558). *การวิเคราะห์ข้อมูลทางสถิติ. เอกสารเผยแพร่ผ่านคอลัมน์
“คลังเอกสารความรู้”. สืบค้นจาก*
<https://www.doa.go.th/share/attachment.php?aid=2730>

จิรภัทร เตชะกุลชัยนันต์. (2555). *การจำแนกประเภทของการวิจัยเพื่อนำไปสู่การเลือกใช้
สถิติ. สืบค้นจาก*
<http://kmcenter.rid.go.th/kcresearch/training/Person27.pdf>

พิมลพรรณ อิศรภักดี. (2557). *แนวคิดการวิจัยทางสังคมศาสตร์และการวิจัยเชิงคุณภาพ.
นครปฐม: สถาบันวิจัยประชากรและสังคม มหาวิทยาลัยมหิดล. สืบค้นจาก*
http://www.rlc.nrct.go.th/ewt_dl.php?nid=1206

วรรณลักษณ์ เมียนเกิด. (2555). *เอกสารบรรยาย โครงการ Research Zone: Phase 67
การวิเคราะห์ข้อมูลเชิงปริมาณและคุณภาพ. สืบค้นจาก*
http://www.rlc.nrct.go.th/ewt_dl.php?nid=1025

- วรวรรณ วงศ์บุตร. (2564). สถิติพื้นฐานสำหรับงานวิจัย (*Basic Statistics*). สืบค้นจาก https://nih.dmsc.moph.go.th/data/data/64/activity/10_8_64/WW_v1.pdf
- สำนักงานคณะกรรมการวิจัยแห่งชาติ. (2561). *คู่มือการตรวจสอบและประเมินผลข้อเสนอการวิจัยของหน่วยงานภาครัฐที่เสนอของบประมาณ ประจำปีงบประมาณ พ.ศ. 2561 ตามมติคณะรัฐมนตรี* [คู่มือออนไลน์]. เผยแพร่เมื่อ 5 สิงหาคม 2559. สืบค้นจาก <https://www.rmutl.ac.th/news/1633>
- สุจรรยา ทรัพย์ศิริโสภ. (2556). เอกสารประกอบการสอนรายวิชาชีวสถิติ บทที่ 1 ความรู้เบื้องต้นทางสถิติ (*Introduction of Statistical*). สืบค้นจาก <https://weatherwing4.6te.net/DataAnalysis%20forWeatherPatterns.pdf>
- สุทิน ชนบุญ. (2560). สถิติและการวิเคราะห์ข้อมูลในงานวิจัยด้านสุขภาพเบื้องต้น. [เอกสารประกอบการบรรยาย]. สืบค้นจาก <https://www.mbuisc.ac.th/woratep/GE4001/2.%20Statistics%20for%20reseach.pdf>
- สุกมาส อังศุโชติ. (2561). *เทคนิคการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร* [เอกสารประกอบการสอน]. มหาวิทยาลัยสุโขทัยธรรมาธิราช. สืบค้นจาก <https://www.stou.ac.th/offices/ore/info/cae/uploads/pdf/636366560441132172.pdf>

2. ภาษาอังกฤษ

(I) Books:

- American Psychological Association. (2020). *Publication manual of the American Psychological Association* (7th ed.). American Psychological Association.
- Bailey, K. D. (1987). *Methods of Social Research*. New York: The Free Press.
- Best, J. W. (1977). *Research in education* (3rd ed.). Prentice-Hall.
- Bivand, R. S., Pebesma, E., & Gómez-Rubio, V. (2013). *Applied spatial data analysis with R* (2nd ed.). Springer.
- Boardman, A. E., Greenberg, D. H., Vining, A. R., & Weimer, D. L. (2018). *Cost-benefit analysis: Concepts and practice* (5th ed.). Cambridge University Press.

- Brynjolfsson, E., & McAfee, A. (2017). *Machine, platform, crowd: Harnessing our digital future*. W. W. Norton & Company.
- Cairo, A. (2016). *The truthful art: Data, charts, and maps for communication*. Berkeley, CA: New Riders.
- Casella, G., & Berger, R. L. (1990). *Statistical inference*. Belmont, CA: Duxbury Press.
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Lawrence Erlbaum Associates.
- Cox, D. R. (2006). *Principles of statistical inference*. Cambridge University Press.
- Creswell, J. W., & Creswell, J. D. (2018). *Research design: Qualitative, quantitative, and mixed methods approaches* (5th ed.). SAGE Publications.
- Cronbach, L. J. (1970). *Essentials of Psychological Testing* (3rd ed.). New York: Harper & Row.
- Dudewicz, E. J. (1988). *Modern mathematical statistics*. New York, NY: Wiley.
- Few, S. (2009). *Now you see it: Simple visualization techniques for quantitative analysis*. Oakland, CA: Analytics Press.
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics* (5th ed.). SAGE Publications.
- Gall, M. D., Borg, W. R., & Gall, J. P. (1996). *Educational Research: An Introduction* (6th ed.). White Plains, NY: Longman.
- Gertler, P. J., Martinez, S., Premand, P., Rawlings, L. B., & Vermeersch, C. M. J. (2016). *Impact evaluation in practice* (2nd ed.). Inter-American Development Bank & World Bank.
- Glass, G. V., & Hopkins, K. D. (1984). *Statistical methods in education and psychology* (2nd ed.). Prentice-Hall.
- GNU PSPP. (2022). *GNU PSPP: Free software for statistical analysis*. Free Software Foundation.
- Gravetter, F. J., & Wallnau, L. B. (2009). *Statistics for the behavioral sciences* (8th ed.). Belmont, CA: Wadsworth.

- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2010). *Multivariate data analysis* (7th ed.). Pearson.
- Howell, D. C. (2007). *Statistical methods for psychology* (6th ed.). Belmont, CA: Thomson Wadsworth.
- Hyndman, R. J., & Athanasopoulos, G. (2021). *Forecasting: Principles and practice* (3rd ed.). OTexts.
- Kerlinger, F. N. (1986). *Foundations of Behavioral Research* (3rd ed.). New York: Holt, Rinehart and Winston.
- Kosslyn, S. M. (2006). *Graph design for the eye and mind*. New York, NY: Oxford University Press.
- Kothari, C. R. (2004). *Research methodology: Methods and techniques* (2nd ed.). New Age International.
- May, T. (1997). *Social Research: Issues, Methods and Process* (2nd ed.). Buckingham: Open University Press.
- Merriam, S. B. (1998). *Qualitative Research and Case Study Applications in Education*. San Francisco: Jossey-Bass.
- Milton, J. S., & Arnold, J. C. (1995). *Introduction to probability and statistics: Principles and applications for engineering and the computer sciences*. Singapore: McGraw-Hill.
- Nachmias, D., & Nachmias, C. (1993). *Research Methods in the Social Sciences* (4th ed.). New York: St. Martin's Press.
- Neuman, W. L. (1991). *Social Research Methods: Qualitative and Quantitative Approaches*. Boston: Allyn and Bacon.
- Norris, P. (2011). *Democratic deficit: Critical citizens revisited*. Cambridge University Press.
- OECD. (2020). *Building capacity for evidence-informed policymaking*. OECD Publishing.
- Shadish, W. R., Cook, T. D., & Campbell, D. T. (2002). *Experimental and quasi-experimental designs for generalized causal inference*. Houghton Mifflin.

- Tufte, E. R. (2001). *The visual display of quantitative information* (2nd ed.). Cheshire, CT: Graphics Press.
- Upton, G., & Cook, I. (2008). *Oxford dictionary of statistics*. Oxford University Press.
- Van Dalen, D. B. (1979). *Understanding Educational Research*. New York: McGraw-Hill.
- Walpole, R. E. (1970). *Introduction to statistics*. New York, NY: Macmillan.
- Ware, C. (2013). *Information visualization: Perception for design* (3rd ed.). Waltham, MA: Morgan Kaufmann.
- Wilkinson, D. (1995). *Research Matters: A Guide to Research in Education and the Social Sciences*. London: Routledge Falmer.
- Yamane, T. (1967). *Statistics: An introductory analysis* (2nd ed.). New York: Harper & Row.
- Yamane, T. (1973). *Statistics: An Introductory Analysis* (3rd ed.). New York: Harper & Row.

(II) Articles:

- Cleveland, W. S., & McGill, R. (1984). Graphical perception: Theory, experimentation and application to the development of graphical methods. *Journal of the American Statistical Association*, 79(387), 531–554.
- Etikan, I., Musa, S. A., & Alkassim, R. S. (2016). Comparison of convenience sampling and purposive sampling. *American Journal of Theoretical and Applied Statistics*, 5(1), 1–4.
- Fung, A. (2015). Putting the public back into governance: The challenges of citizen participation and its future. *Public Administration Review*, 75(4), 513–522.
- Janssen, M., Charalabidis, Y., & Zuiderwijk, A. (2017). Benefits, adoption barriers and myths of open data and open government. *Information Systems Management*, 34(3), 229–240.
- Krejcie, R. V., & Morgan, D. W. (1970). Determining sample size for research activities. *Educational and Psychological Measurement*, 30(3), 607–610.

- Likert, R. (1932). *A technique for the measurement of attitudes*. Archives of Psychology, 22(140), 1–55.
- Love, J., Selker, R., Marsman, M., Jamil, T., Dropmann, D., Verhagen, J., & Wagenmakers, E. J. (2019). JASP: Graphical statistical software for common statistical designs. *Journal of Statistical Software*, 88(2), 1–17.
- Sirin, S. R. (2005). Socioeconomic status and academic achievement: A meta-analytic review of research. *Review of Educational Research*, 75(3), 417–453.

(II) Online:

- Chappelow, J. (2025). *Statistics: Definition, Types and Importance*. Retrieved from <https://www.investopedia.com/terms/s/statistics.asp>
- Frost, J. (2021). Statistical inference. *Statistics by Jim*. Retrieved from <https://statisticsbyjim.com/hypothesis-testing/statistical-inference/>
- Penn State University. (2023). Statistical inference and estimation. *STAT 504: Analysis of Discrete Data*. Retrieved from <https://online.stat.psu.edu/stat504/lesson/statistical-inference-and-estimation>
- R Core Team. (2023). *R: A language and environment for statistical computing* [Computer software]. R Foundation for Statistical Computing. Retrieved from <https://www.R-project.org/>



ประวัติผู้แต่ง

- ชื่อ-นามสกุล : ผู้ช่วยศาสตราจารย์ ดร.สุภัทรชัย สีสะใบ
- วัน เดือน ปีเกิด : วันที่ 16 เมษายน พ.ศ. 2529
- ภูมิลำเนาที่เกิด : 13 หมู่ 2 ตำบลห้วยแย้ อำเภอนองบัวระเหว จังหวัดชัยภูมิ
- การศึกษา : พ.ศ. 2553 หลักสูตรพุทธศาสตรบัณฑิต วิชาเอกภาษาอังกฤษ
คณะมนุษยศาสตร์ มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย
วิทยาเขตนครราชสีมา
พ.ศ. 2555 หลักสูตรพุทธศาสตรมหาบัณฑิต
สาขาวิชารัฐประศาสนศาสตร์ ภาควิชารัฐศาสตร์
คณะสังคมศาสตร์ มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย
พ.ศ. 2563 หลักสูตรปรัชญาดุษฎีบัณฑิต
สาขาวิชารัฐประศาสนศาสตร์ ภาควิชารัฐศาสตร์
คณะสังคมศาสตร์ มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย
- ประสบการณ์การทำงาน : อาจารย์ประจำภาควิชารัฐศาสตร์ คณะสังคมศาสตร์
มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย
- สังกัด : หลักสูตรบัณฑิตศึกษา ภาควิชารัฐศาสตร์ คณะสังคมศาสตร์
มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย
- ตำแหน่ง : อาจารย์
- ที่อยู่ปัจจุบัน : 9/6 หมู่ 3 ตำบลขยาย อำเภอบางปะหัน จังหวัดพระนครศรีอยุธยา



สถิติและการวิจัยทางสังคมศาสตร์:

แนวคิด วิธีการ และการประยุกต์

ผศ.ดร.สุภัทรชัย ลีสะไบ

ภาควิชารัฐศาสตร์ คณะสังคมศาสตร์ มหาวิทยาลัยมหาจุฬาลงกรณราชวิทยาลัย