



มหาวิทยาลัยมหิดล
Mahidol University
"Wisdom of The Land"



การเก็บรวบรวมและวิเคราะห์ผลข้อมูล R2R (Data Collection and Analysis for R2R)

โดย

รองศาสตราจารย์ ดร.สมบูรณ์ ศิริสรหิรัญและทีมงาน
คณะสังคมศาสตร์และมนุษยศาสตร์ มหาวิทยาลัยมหิดล
โทร.098-9165624 somboon.sir@mahidol.ac.th



การกำหนดประชากรและการเลือกตัวอย่าง

- การวิจัยเรื่องหนึ่งๆ ต้องเก็บรวบรวมข้อมูลทั้งหมดจากกลุ่มประชากรเป้าหมาย (Target Population) มาศึกษา
- แต่โดยทั่วไปการวิจัยไม่สามารถเก็บรวบรวมข้อมูลจากทุกหน่วยได้ เนื่องจากเสียเวลา เปลืองค่าใช้จ่าย ใช้กำลังคนจำนวนมาก จึงนิยมใช้วิธีการหาข้อมูลจากตัวแทน (Representative) ของประชากร ซึ่งเรียกว่า กลุ่มตัวอย่าง (Samples) แทน



การกำหนดประชากรและการเลือกตัวอย่าง (ต่อ)

และศึกษาดูข้อมูลเหล่านั้นว่าเป็นตัวแทนที่ดีของประชากรหรือไม่ การกำหนดขอบเขตของประชากรและกลุ่มตัวอย่างมีผลต่อวิธีการเก็บรวบรวมข้อมูล การสร้างเครื่องมือวิจัย การวิเคราะห์ข้อมูล การสุ่มตัวอย่างที่ไม่เหมาะสมทำให้ผลวิจัยคลาดเคลื่อนได้ เนื่องจากกลุ่มตัวอย่างไม่ได้เป็นตัวแทนที่แท้จริงของประชากร สิ่งที่ต้องพิจารณาคือ จะใช้วิธีการสุ่มตัวอย่างแบบใด และจำนวนกลุ่มตัวอย่างมากน้อยแค่ไหนจึงจะถือว่าเหมาะสม



การคัดเลือกกลุ่มตัวอย่าง

- การสุ่มแบบง่าย (simple random sampling)
- การสุ่มแบบแบ่งชั้น (stratified random sampling)
- การสุ่มแบบแบ่งกลุ่ม (cluster random sampling)
- การสุ่มแบบมีระบบ (systematic random sampling)
- การสุ่มหลายขั้นตอน (multi-stage random sampling)



การเก็บรวบรวมข้อมูล

- ผลวิจัยจะมีความเที่ยงตรง (validity) เชื่อถือได้ (reliability) ได้มากน้อยเพียงใด ขึ้นอยู่กับข้อมูลที่เก็บรวบรวม
- การเก็บรวบรวมข้อมูลมีหลายวิธีขึ้นอยู่กับประเภทของข้อมูล และต้องเก็บให้ตรงตามวัตถุประสงค์ และสมมติฐาน อีกทั้งคำนึงถึงการนำข้อมูลไปวิเคราะห์
- ข้อพิจารณาเกี่ยวกับการเก็บรวบรวมข้อมูล ได้แก่ แหล่งข้อมูล ประเภทของข้อมูล ลักษณะของข้อมูล และเครื่องมือที่ใช้ในการเก็บรวบรวมข้อมูล



การเก็บรวบรวมข้อมูลสนาม

ที่สำคัญและใช้กันทั่วไปมี 3 วิธี คือ

- การสังเกต (observation)
- การสัมภาษณ์ (interview)
- การใช้แบบสอบถาม (questionnaire)
 - คำถามปลายเปิด
 - คำถามปลายปิด



การวิเคราะห์ข้อมูล (Data Analysis)

- นำข้อมูลที่เก็บมาประมวลผล เป็นการจัดระเบียบของข้อมูลใหม่ เพื่อให้สามารถนำไปวิเคราะห์หาคำตอบตามวัตถุประสงค์ที่กำหนดได้
- การประมวลผลจะสลับซับซ้อนมากขึ้นขึ้นอยู่กับขนาดของกลุ่มตัวอย่าง
- สถิติที่ใช้ในการวิเคราะห์ข้อมูล คือ สถิติอ้างอิง (Inferential Statistics) เป็นการนำค่าสถิติเชิงพรรณนาที่ได้ในรูปของการประมวลผลมาสรุปอ้างอิงถึงกลุ่มประชากรที่ศึกษา เพื่อสรุปผลและตอบปัญหาการวิจัย



การวิเคราะห์ข้อมูล

การวิเคราะห์ข้อมูลแบ่งได้ 2 ลักษณะ

1. การวิเคราะห์ข้อมูลด้วยเหตุผล
2. การวิเคราะห์ข้อมูลด้วยวิธีการทางสถิติ

การวิจัยเชิงคุณภาพ

การวิเคราะห์และสังเคราะห์ข้อมูล เพื่อให้ได้ความหมายของปรากฏการณ์ที่เกิดขึ้นโดยการเชื่อมโยงความสัมพันธ์ระหว่างปรากฏการณ์ที่เกิดขึ้นกับสภาพแวดล้อมทั้งหมดและหาข้อสรุปของโครงสร้างของปรากฏการณ์ในเชิงอุปนัย

การวิจัยเชิงปริมาณ
ต้องอาศัยวิธีการทางสถิติมาพรรณนาหรือทดสอบสมมติฐานที่ตั้งไว้



กระบวนการจัดการข้อมูล

- การจัดแบ่งกลุ่มข้อมูล
- การจัดเรียงข้อมูล
- การสรุปผล
- การคำนวณ
 - Coding → ลงข้อมูลในคอมพิวเตอร์ → Editing
 - ความแตกต่าง
 - ข้อมูลจำแนกนับ เลขจำนวนเต็ม Chi Square
 - ข้อมูลเชิงปริมาณ มีทศนิยม T test
 - ความสัมพันธ์ Regress



การกำหนดขนาดตัวอย่าง

**เชิงปริมาณ ต้องเพียงพอเป็นตัวแทนที่ดี
ของประชากรได้ ใช้เทคนิคในการเลือก
อย่างน่าเชื่อถือ**

**เชิงคุณภาพ ต้องตรงตามเงื่อนไข เป็น
บุคคลสำคัญที่ให้ข้อมูลได้ตรงตามที่
ต้องการ**

***ขึ้นกับ เวลา งบประมาณ และขนาดของ R2R**



การคำนวณขนาดตัวอย่างด้วยสูตร

$$n = \left(\frac{Z_{\alpha} \sigma}{e} \right)^2$$

n คือ ขนาดของกลุ่มตัวอย่าง

Z คือ ค่ามาตรฐานตามระดับความมีนัยสำคัญ

σ คือ ค่าประมาณของส่วนเบี่ยงมาตรฐานของประชากร

e คือ ความคลาดเคลื่อนที่ยอมให้เกิดขึ้น (ค่า $\mu - \bar{X}$)



การคำนวณขนาดตัวอย่างด้วยสูตร Taro Yamane

$$n = \frac{N}{1 + Ne^2}$$

n = ขนาดของกลุ่มตัวอย่าง

N = ขนาดของประชากร

e = ความคลาดเคลื่อนของข้อมูล

(Taro Yamane, 1976)

***การคำนวณขนาดตัวอย่างด้วยสูตร
ใน R2R อาจใช้ตารางสำเร็จรูปก็อาจ
ทำให้หาขนาดตัวอย่างได้ง่ายขึ้น**



ตาราง Taro Yamane

จำนวน ประชากร	ขนาดของตัวอย่างประชากรสำหรับความคลาดเคลื่อนที่กำหนด (คิดเป็นร้อยละ)					
	± 1	± 2	± 3	± 4	± 5	± 10
500	b*	b	b	b	222	83
1,000	b	b	b	385	286	91
1,500	b	b	638	441	316	94
2,000	b	b	714	476	333	95
2,500	b	1,250	769	500	345	96
3,000	b	1,364	811	517	353	97
3,500	b	1,458	843	530	359	97
4,000	b	1,538	870	541	364	97
4,500	b	1,607	891	549	367	98
5,000	b	1,667	909	556	370	98
6,000	b	1,765	938	566	375	98
7,000	b	1,842	959	574	378	99
8,000	b	1,905	976	580	381	99
9,000	b	1,957	989	584	383	99
10,000	5,000	2,000	1,000	588	385	99
15,000	6,000	2,143	1,034	600	390	99
20,000	6,667	2,222	1,053	606	392	100
25,000	7,143	2,273	1,064	610	394	100
50,000	8,333	2,381	1,087	617	397	100
100,000	9,091	2,439	1,099	621	398	100
$\rightarrow \infty$	10,000	2,500	1,111	625	400	100

b* เป็นค่าที่ไม่สามารถคำนวณได้จากสูตร



ตาราง Krejcie and Morgan

ประชากร	ขนาดของกลุ่ม ตัวอย่าง	ประชากร	ขนาดของกลุ่ม ตัวอย่าง	ประชากร	ขนาดของกลุ่ม ตัวอย่าง
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
210	136	1100	285	100000	384

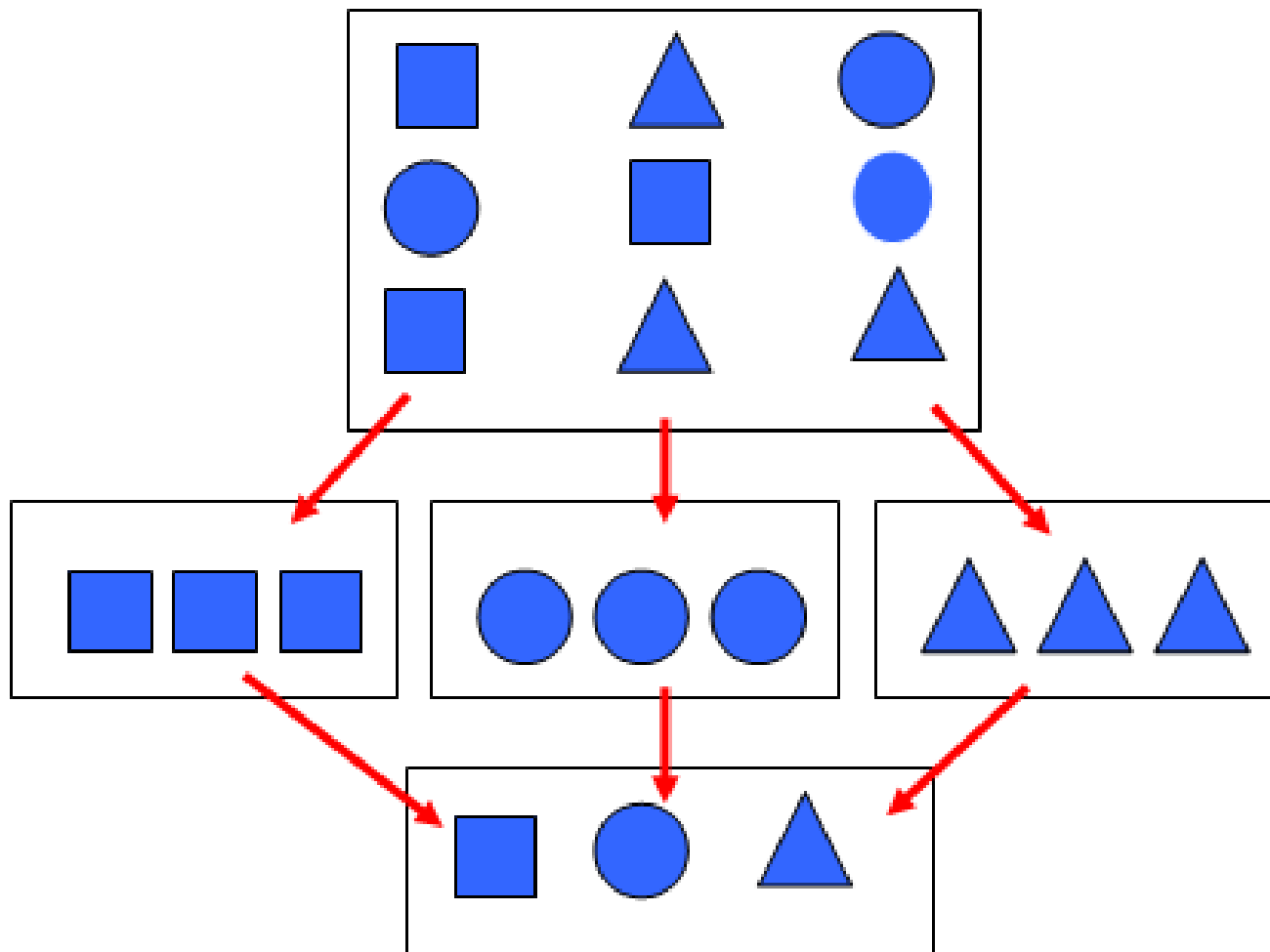


เทคนิคการสุ่มตัวอย่าง

- การสุ่มตัวอย่างแบบอิงค่าความน่าจะเป็น (Probability Sampling)
 - การสุ่มตัวอย่างแบบอย่างง่าย
 - การสุ่มตัวอย่างแบบมีระบบ
 - การสุ่มตัวอย่างแบบแบ่งชั้น
 - การสุ่มตัวอย่างแบบแบ่งกลุ่ม
 - การสุ่มตัวอย่างแบบหลายขั้นตอน
- การสุ่มตัวอย่างแบบไม่อิงค่าความน่าจะเป็น (Nonprobability Sampling)
 - การสุ่มตัวอย่างแบบบังเอิญ
 - การสุ่มตัวอย่างแบบกำหนดจำนวน (สัดส่วน / โควตา)
 - การสุ่มตัวอย่างแบบเจาะจง

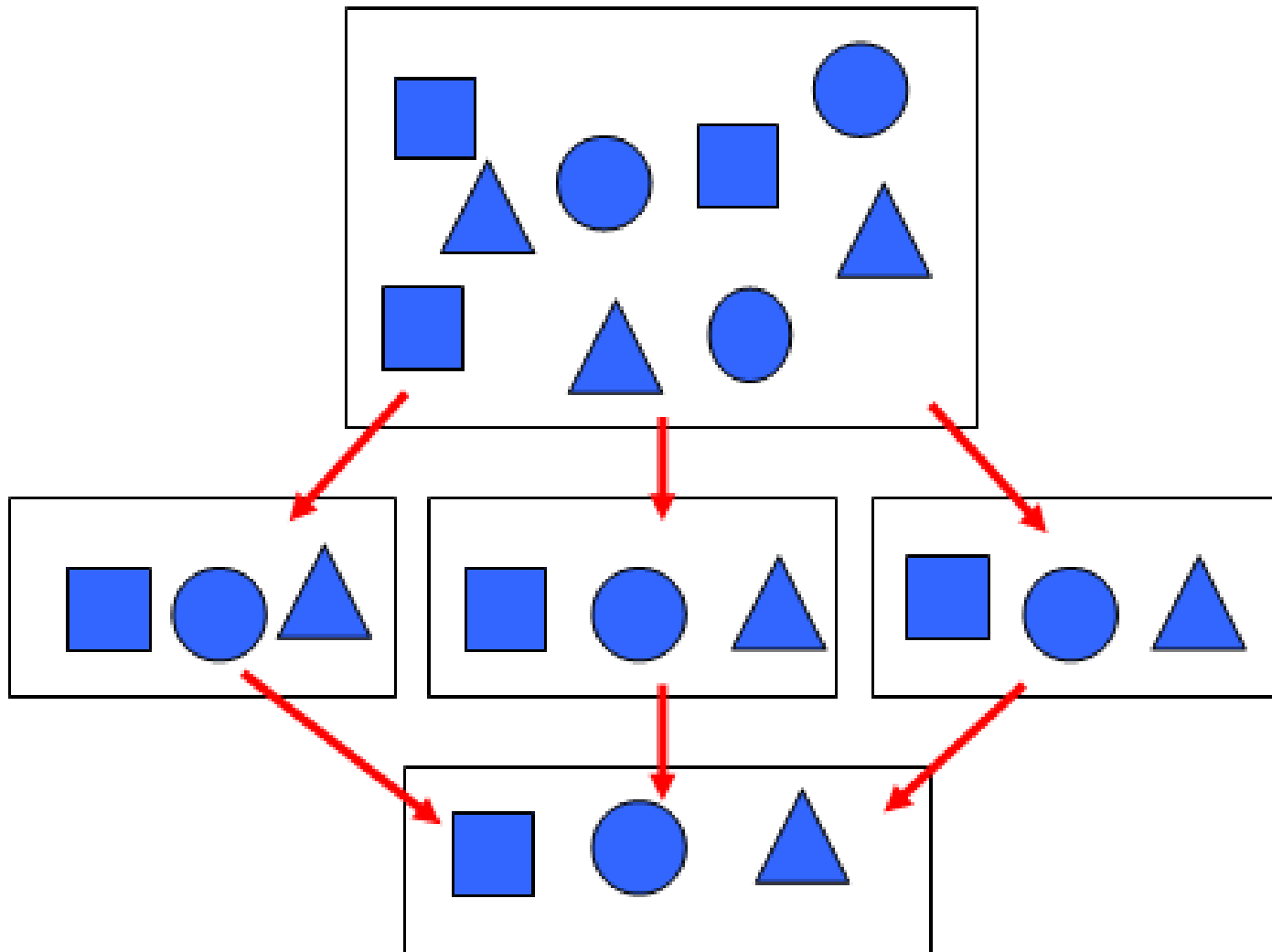


การสุ่มแบบแบ่งชั้น





การสุ่มแบบแบ่งกลุ่ม



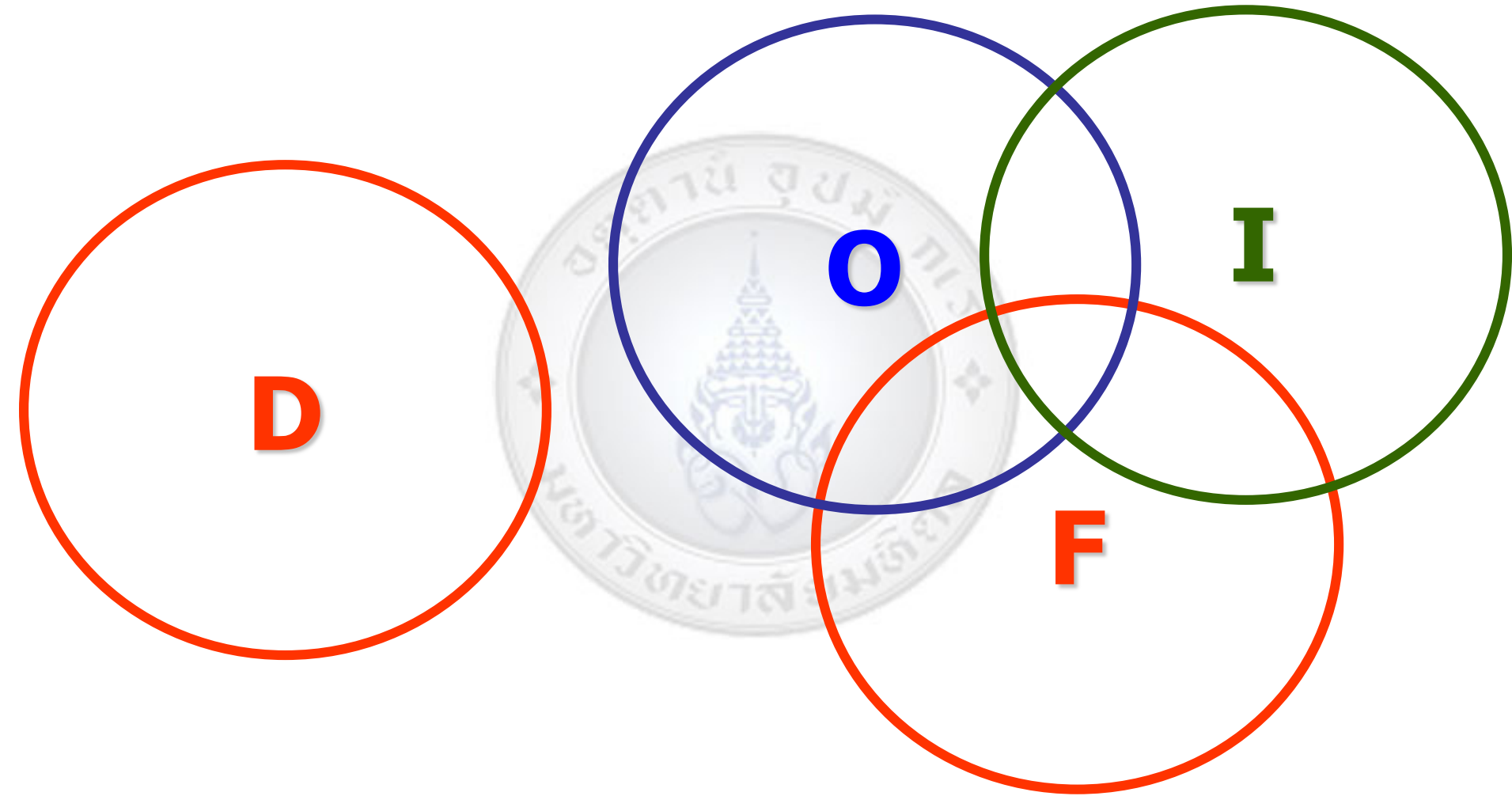


การวิจัยเชิงคุณภาพ

- การเก็บรวบรวมข้อมูล ที่สำคัญ ใน R2R
 - การวิเคราะห์เนื้อหาของเอกสาร (Document Search)
 - การสังเกต (Observation) การสังเกตในวิจัยเชิงคุณภาพมี 2 แบบ คือ
 - การสังเกตแบบมีส่วนร่วม (participation observation)
 - การสังเกตแบบไม่มีส่วนร่วม (non-participation observation)
 - การสัมภาษณ์ (Interview) การสัมภาษณ์เป็นการเจาะลึกประเด็นต่าง ๆ ที่ผู้วิจัยสนใจ อาจใช้สัมภาษณ์เป็นรายบุคคลหรือเป็นกลุ่มก็ได้ มีหลายประเภท อาจแบ่งได้ดังนี้
 - การสัมภาษณ์แบบเป็นทางการ (formal interview)
 - การสัมภาษณ์แบบไม่เป็นทางการ (informal interview)
 - การประชุมกลุ่ม (Focus Group) การประชุมเพื่อเก็บข้อมูลในการศึกษา เป็นรูปแบบสำคัญแบบหนึ่งโดยใช้กระบวนการกลุ่ม



Main 4 Tools





การหาคุณภาพของเครื่องมือสำหรับวิจัย

มีกระบวนการหาคุณภาพเครื่องมือเพื่อความ
น่าเชื่อถือของงานวิจัย ด้วยวิธีต่างๆ เช่น

ความเที่ยงตรง (Validity)

ความเชื่อมั่น (Reliability)

ดัชนีความยากของข้อสอบหรือดัชนีค่าความง่ายของคำถามวัดความรู้

ดัชนีค่าอำนาจจำแนกของคำถามวัดความรู้

ความเป็นปรนัย (Objectivity) ของคำถามวัดความรู้

ความสะดวกใช้ (Usability)



การหาคุณภาพของเครื่องมือสำหรับ R2R

ความเที่ยงตรง (Validity)

เป็นคุณภาพของแบบทดสอบที่หมายถึง
แบบทดสอบที่สามารถวัดได้ตรงตามลักษณะ
หรือจุดประสงค์ที่ต้องการจะวัด

ความเชื่อมั่น (Reliability)

ความคงที่ของคะแนนที่ได้จากการสอบ
นักเรียนคนเดียวกันหลายครั้งในแบบทดสอบ
ชุดเดิม



ความเที่ยงตรง (Validity)

ความเที่ยงตรงเชิงเนื้อหา ความเที่ยงตรงตาม
โครงสร้าง ความเที่ยงตรงตามสภาพ และความ
เที่ยงตรงเชิงพยากรณ์

ความเที่ยงตรงเชิงเนื้อหา (Content Validity) เป็นความเที่ยงตรงที่หาโดยการให้
ผู้เชี่ยวชาญพิจารณาว่าข้อสอบ หรือ ข้อคำถามแต่ละข้อ วัดได้ตรงตามสิ่งที่ต้องการวัด
เนื้อหาหรือวัตถุประสงค์การเรียนรู้มากน้อยเพียงใด โดยใช้เกณฑ์การประเมิน ดังนี้

ให้คะแนน +1	หมายถึง	แน่ใจว่าข้อสอบวัดจุดประสงค์/เนื้อหานั้น
ให้คะแนน 0	หมายถึง	ไม่แน่ใจว่าข้อสอบวัดจุดประสงค์/เนื้อหานั้น
ให้คะแนน -1	หมายถึง	แน่ใจว่าข้อสอบไม่วัดจุดประสงค์/เนื้อหานั้น

แล้วนำข้อมูลที่ได้จากการพิจารณาของผู้เชี่ยวชาญ หาค่า
ความสอดคล้องระหว่างข้อคำถามแต่ละข้อกับจุดประสงค์
หรือเนื้อหา (Index of Item-Objective Congruence หรือ IOC)



ค่า IOC

$$\text{IOC} = \Sigma R / N$$

เมื่อ ΣR แทน ผลรวมของคะแนนการ
พิจารณาของผู้เชี่ยวชาญ
 N แทน จำนวนผู้เชี่ยวชาญ

เกณฑ์การตัดสินค่า IOC ถ้ามีค่า 0.50 ขึ้น
ไป แสดงว่า ข้อคำถามนั้นวัดได้ตรง
จุดประสงค์ หรือตรงตามเนื้อหานั้น แสดงว่า
ข้อคำถามข้อนั้นใช้ได้



ตัวอย่างค่า IOC

ตัวอย่าง การคำนวณและการแปลผลค่า IOC

ข้อ (รายการ)	คะแนนความเห็นของผู้เชี่ยวชาญ					รวม	ค่า IOC	สรุปผล
	คนที่ 1	คนที่ 2	คนที่ 3	คนที่ 4	คนที่ 5			
1	+1	+1	+1	+1	+1	5	1	ใช้ได้
2	0	+1	+1	0	+1	3	0.6	ใช้ได้
3	+1	0	-1	0	0			
4	+1	+1	-1	+1	+1			
5	0	0	-1	0	-1			



ความเชื่อมั่น (Reliability)

ความเชื่อมั่นเป็นดัชนีหรือระดับที่ชี้ให้ทราบว่า แบบทดสอบมีความคงที่ในการวัดถึงแม้ว่าจะวัดก็ครั้งก็ตาม

การหาความเชื่อมั่นของแบบทดสอบ มีหลายวิธี เช่น

- การสอบซ้ำ (Test-retest)
- การใช้แบบทดสอบคู่ขนาน (Parallel forms)
- การแบ่งครึ่งแบบทดสอบ (Split-half)
- การใช้วิธีการของ คูเดอร์ ริชาร์ดสัน (Kuder-Richardson procedures)
- การใช้สูตรของ ครอนบัค (Cronbach) คือ α -Coefficient



-ถ้าการให้คะแนนเป็น 0 และ 1 กล่าวคือ ทำถูกต้อง 1 คะแนน ทำผิดได้ 0 คะแนน จะใช้วิธีการของคูเดอร์ - ริชาร์ดสัน (Kuder Richardson:KR-20, KR-21)

-ถ้าคะแนนไม่ใช่ 0 กับ 1 คือ ให้คะแนนลักษณะใดก็ได้ เช่น แบบสอบถาม หรือ แบบวัดเจตคติ ให้คะแนนเป็น 5, 4, 3, 2 และ 1 หรือ 3, 2 และ 1 แบบทดสอบแบบอัตนัย ที่ให้คะแนนแต่ละข้อต่างกัน เช่น 5 หรือ 8 หรือ 10 คะแนน จะใช้สูตรของ ครอนบัค (Cronbach) คือ α - Coefficient



ค่าความเชื่อมั่นมีค่าตั้งแต่ 0 ถึง 1.00

การพิจารณาค่าความเชื่อมั่น ไม่มี
กฎเกณฑ์ที่ตายตัว แต่มีค่ามากก็ยิ่งดี
กล่าวคือ ถ้าแบบทดสอบใดมีความ
เชื่อมั่นสูง จะทำให้ผู้วิจัยมีความมั่นใจใน
เครื่องมือที่ใช้เก็บรวบรวมข้อมูลยิ่งขึ้น
ซึ่งส่งผลต่อความเชื่อมั่นในผลการวิจัย

ค่าความเชื่อมั่นจาก Output Windows ของโปรแกรม SPSS

SPSS Statistics Data Editor window showing a dataset with 14 rows and 3 columns (เลขที่, เพศ, and an unlabeled column). The 'Analyze' menu is open, and 'Scale' is selected, leading to the 'Reliability Analysis...' dialog box.

	เลขที่	เพศ
1	1	2
2	2	2
3	3	1
4	4	1
5	5	2
6	6	2
7	7	1
8	8	2
9	9	1
10	10	2
11	11	2
12	12	2
13	13	2
14	14	2

Reliability Analysis...

ที่เมนู Analyze เลือกทำห้ถึง
Scale และเลือกตัวเลือก Reliability
Analysis เพื่อคำนวณหาค่าความ
เชื่อมั่นในข้อนี้

Scale: ALL VARIABLES

Case Processing Summary

		N	%
Cases	Valid	30	100.0
	Excluded ^a	0	.0
	Total	30	100.0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
.858	3

ค่าความเชื่อมั่นของแบบสอบถามทั้งชุด หลังจากที่ได้ตัดข้อคำถาม "ถูกต้อง" ออกแล้ว จะมีค่าเพิ่มขึ้น

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted
พอใจในพิธีการผู้จัด	6.80	5.614	.795	.749
พอใจผู้โทรเข้าร่วม	6.77	5.584	.689	.843
พอใจความชัดเจน	6.57	5.357	.721	.814



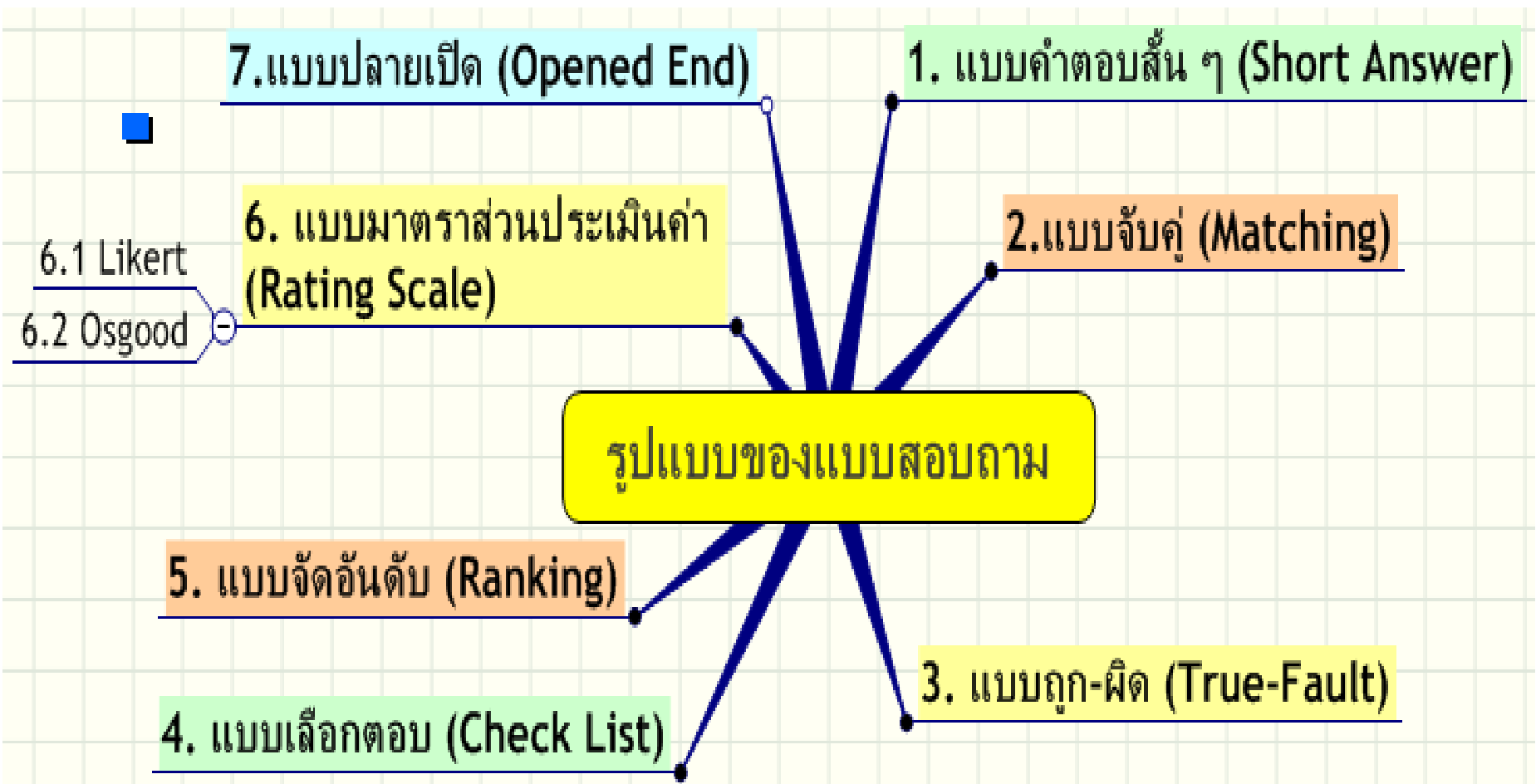
การวิเคราะห์ข้อมูลเชิงคุณภาพ

- การวิเคราะห์ข้อมูลโดยการตีความสร้างข้อสรุปแบบอุปนัย (inductive)
 - ได้จากการสังเกตและการสัมภาษณ์ที่ได้จดบันทึกไว้จากสิ่งที่เป็นรูปธรรมหรือปรากฏการณ์ที่มองเห็น
 - ผู้วิจัยได้เห็นหลาย ๆ เหตุการณ์และได้ทำการตรวจสอบข้อมูลแบบสามเส้า
 - เขียนเป็นประโยคหรือข้อความเพื่อสร้างข้อสรุปได้ตามกรอบแนวคิดทฤษฎีหรือเพื่อตอบปัญหาของการวิจัย ข้อมูลที่ไม่ต้องการจะถูกกำจัดออกไปได้
- การวิเคราะห์ข้อมูลโดยการวิเคราะห์เนื้อหา (content analysis)
 - ได้จากการศึกษาเอกสาร (Document Research)
 - ผู้วิจัยต้องคำนึงถึงบริบท (context) หรือสภาพแวดล้อมของข้อมูลเอกสารที่นำมาวิเคราะห์ประกอบด้วยว่ามีการเปลี่ยนแปลงไปอย่างไร

การวิเคราะห์ข้อมูลเชิงคุณภาพเป็นข้อความแบบบรรยาย (descriptive) ไม่มีสูตรสำเร็จตายตัว ขึ้นอยู่กับประเด็นหรือปัญหาที่จะวิเคราะห์และการเลือกของนักวิจัย ดังนั้นการมีกรอบความคิดหรือทฤษฎีที่หลากหลายจะมีความสำคัญอย่างยิ่งในการช่วยวิเคราะห์ข้อมูลได้ลึกซึ้งและสร้างข้อสรุปที่หนักแน่น



แบบสอบถาม ใน R2R เชิงปริมาณ





Likert scale

แบบลิเคิร์ต นิยมใช้มากกับแบบสอบถามที่วัดเจตคติ ความเชื่อ ความพึงพอใจ ความเห็นต่อเรื่องต่างๆ โดยกำหนดรูปแบบออกเป็นระดับความคิดเห็นของผู้ตอบ 1-5 ระดับ เพื่อเป็นการบอกขนาดความเข้มของความคิดเห็นจากผู้จะตอบใน R2R โดยข้อคำถามจะประกอบด้วยข้อคำถามททั้งทางบวกและทางลบ ตัวอย่างเช่น



ตัวอย่าง Likert scale

ข้อคำถาม	ไม่เห็นด้วย อย่างยิ่ง	ไม่เห็น ด้วย	ไม่แน่ใจ หรือ เฉย ๆ	เห็นด้วย	เห็นด้วย อย่างยิ่ง
ฉันคิดว่าการมาฝาก ครรภ์ตามกำหนดนัดมี ความสำคัญมาก					
ในระยะตั้งครรภ์ การ ปฏิบัติตามที่แพทย์สั่ง นั้นยุ่งยากมาก					



Osgood

วิธีการหาความแตกต่างของความหมาย

สถานีอนามัยในชุมชนของฉัน

ดี	7	6	5	4	3	2	1	เลว
บริการ เร็ว	7	6	5	4	3	2	1	บริการช้า
ใหญ่	7	6	5	4	3	2	1	เล็ก
สะอาด	7	6	5	4	3	2	1	สกปรก
ใหม่	7	6	5	4	3	2	1	เก่า



ระดับการวัดข้อมูล

ลักษณะข้อมูลและระดับการวัด (Data character and Level of measurement)

1. ข้อมูลระดับนามบัญญัติ (Nominal data)
2. ข้อมูลระดับเรียงลำดับ (Ordinal data)
3. ข้อมูลระดับอันตรภาค (Interval data)
4. ข้อมูลระดับอัตราส่วน (Ratio data)



การวิเคราะห์ข้อมูลเชิงปริมาณ

สถิติเชิงพรรณนา (Descriptive Statistics)

เป็นระเบียบวิธีทางสถิติ ที่มุ่งอธิบายให้เห็นภาพของข้อมูลทั้งหมด ได้แก่ การนำเสนอข้อมูลในรูปตาราง แผนภูมิ กราฟ หรือรูปภาพต่าง ๆ การแปลความหมายข้อมูล การคำนวณ และการตีความหมายของค่ากลางต่าง ๆ เป็นต้น

สถิติที่ใช้ในการบรรยายหรือพรรณนาคุณลักษณะของข้อมูล ได้แก่ ความถี่ ร้อยละ ค่าเฉลี่ย มัธยฐาน ส่วนเบี่ยงเบนมาตรฐาน เป็นต้น



สถิติเชิงพรรณนา

(Descriptive statistics)

- การแจกแจงความถี่ - จัดทำตารางแจกแจงความถี่
- การหาค่าสัดส่วน/ ร้อยละ - คำนวณค่า %
- การวัดแนวโน้มเข้าสู่ส่วนกลาง - หาค่าเฉลี่ย
- การหาค่าการกระจาย - หาค่า ส่วนเบี่ยงเบนมาตรฐาน (SD)



การนำเสนอและการเขียนรายงาน

1. การนำเสนอข้อมูลอย่างไม่เป็นแบบแผน

1.1 บทความ

1.2 บทความกึ่งตาราง

2. การนำเสนอข้อมูลอย่างเป็นแบบแผน

2.1 ตาราง

2.2 แผนภูมิ

2.3 กราฟ

จากนั้นทำเป็นรูปเล่มรายงานวิจัย



การนำเสนอข้อมูลหรือการแสดงผล การวิเคราะห์ข้อมูล

ตารางตัวแปรเดียว

ตารางที่ 1 สถานภาพทั่วไปของครูผู้สอนที่เป็นกลุ่มตัวอย่างในการวิจัยเรื่องความต้องการในการจัดหาหลักสูตรท้องถิ่น

สถานภาพทั่วไป	จำนวน	ร้อยละ
เพศ		
- ชาย	4	25.0
- หญิง	12	75.0
ระดับการศึกษา		
- ปริญญาตรี	12	75.0
- ปริญญาโท	4	25.0
ประสบการณ์ในการสอน		
- 1-5 ปี	8	50.0
- 6 - 10 ปี	3	18.7
- 10 ปีขึ้นไป	5	31.3
การได้รับความรู้เกี่ยวกับการพัฒนาหลักสูตรท้องถิ่น		
- ไม่เคย	1	6.2
- เคย	15	93.8



ตัวอย่างผลการประเมินโครงการฝึกอบรม จากแบบสำรวจความคิดเห็น

ส่วนที่ 1 ความเหมาะสมของเนื้อหาหลักสูตรและรูปแบบการฝึกอบรม (n=48)

เนื้อหา/รูปแบบการฝึกอบรม	ระดับความพึงพอใจ (ร้อยละ)			
	มากที่สุด	มาก	น้อย	น้อยที่สุด
1. ความครบถ้วนของหลักสูตร	65.5	34.5	-	-
2. การจัดลำดับเนื้อหา/ความต่อเนื่อง	75.9	24.1	-	-
3. รูปแบบการอบรม	69.0	24.1	6.9	-
4. ระยะเวลาการอบรม	51.7	37.9	10.3	-
5. ช่วงเวลาการอบรม	51.7	41.4	6.9	-
6. เอกสารประกอบการอบรม	58.6	41.4	-	-
โดยภาพรวม	72.4	27.6	-	-



ตารางสองตัวแปรหรือตารางที่มีตัวแปร 2 ตัว

ตารางที่2 จำนวนผู้ป่วยแยกตามที่ทำงานแบ่งตามขนาดโรงงานและที่ตั้ง

ขนาดโรงงาน	แหล่งที่ตั้ง	
	กรุงเทพฯ	ต่างจังหวัด
ขนาดใหญ่	135	184
ขนาดกลาง	1,025	3,250
ขนาดเล็ก	3,260	34,580
รวม	4,420	38,014



ตาราง3ตัวแปรหรือตารางที่มีตัวแปร3ตัว

ตารางที่3 จำนวนผู้ป่วยที่ศึกษาแยกตามขนาด โรงงาน แหล่งที่ตั้ง และประเภทของกิจการ

ขนาดโรงงาน	แหล่งที่ตั้ง			
	กรุงเทพฯ		ต่างจังหวัด	
	ต่างชาติ	ไทย	ต่างชาติ	ไทย
ขนาดใหญ่	80	55	96	88
ขนาดกลาง	190	835	350	2,900
ขนาดเล็ก	20	3,240	-	34,580
รวม	290	4,130	446	37,568



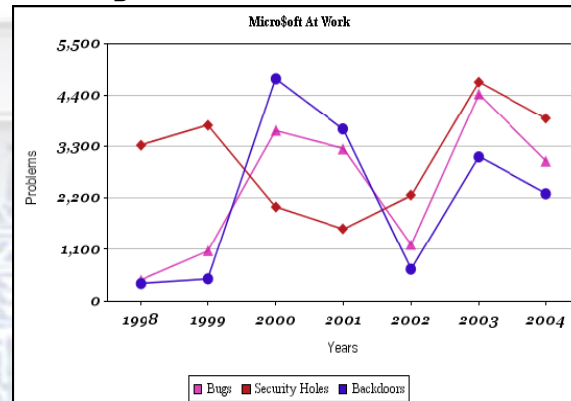
การนำเสนอผลการแจกแจงความถี่

• การนำเสนอด้วยแผนภูมิ

แผนภูมิภาพ (pictograph)



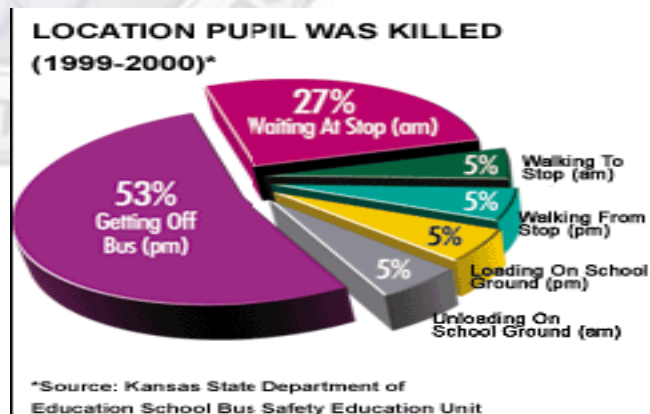
แผนภูมิเส้น (Trend charts)



แผนภูมิแท่ง (Bar charts)



แผนภูมิวง (Pie diagram)





การวัดแนวโน้มเข้าสู่ส่วนกลาง

(Measures of central tendency, Measure of location)

- **ค่าเฉลี่ย (Mean)** เป็นค่าวัดแนวโน้มเข้าสู่ส่วนกลางที่นิยมใช้มากที่สุด เป็นค่าที่เกิดจากการนำเอาค่าของหน่วยข้อมูลทุกๆ หน่วยที่เก็บรวบรวมได้มาบวกกัน แล้วหารด้วยจำนวนหน่วย ข้อมูลทั้งหมด

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



เราจะวัดค่าจาก Likert Scale ได้อย่างไร

ข้อคำถาม	ไม่เห็นด้วย อย่างยิ่ง	ไม่เห็น ด้วย	ไม่แน่ใจ หรือ เฉย ๆ	เห็นด้วย	เห็นด้วย อย่างยิ่ง
ฉันคิดว่าการมาฝาก ครรภ์ตามกำหนดนัดมี ความสำคัญมาก					
ในระยะตั้งครรภ์ การ ปฏิบัติตามที่แพทย์สั่ง นั้นยุ่งยากมาก					

***Likert Scale จะเป็น Ordinal Data หา Mean ไม่ได้**



ต้องปรับ Likert Scale ให้เป็นค่าตัวเลข

ข้อคำถาม	ไม่เห็นด้วยอย่างยิ่ง	ไม่เห็นด้วย	ไม่แน่ใจ หรือ เฉย ๆ	เห็นด้วย	เห็นด้วยอย่างยิ่ง
ต้องกำหนดค่าเป็นตัวเลขให้ ทุกค่าคำตอบ	5	4	3	2	1
กำหนดค่าการแปลผล	... - 5.00	... - ...	... - ...	... - ...	1 - ...



แปลผลค่าเฉลี่ย

(Max-Min)/ จำนวนชั้น

เนื้อหา/รูปแบบการฝึกอบรม	ระดับความพึงพอใจ (ร้อยละ)			
	มากที่สุด	มาก	น้อย	น้อยที่สุด
	4	3	2	1

จำนวนชั้น	4
Max	4
Min	1



แปลผลค่าเฉลี่ย

จำนวนชั้น	4
Max	4
Min	1

$(\text{Max}-\text{Min}) / \text{จำนวนชั้น}$

$$(4-1) / 4 = 3/4 = 0.75$$



แปลผลค่าเฉลี่ย

เนื้อหา/รูปแบบการฝึกอบรม	ระดับความพึงพอใจ (ร้อยละ)			
	มากที่สุด	มาก	น้อย	น้อยที่สุด
ค่าจริง	4	3	2	1
ค่าแปลผล	3.26 - 4	2.51 – 3.25	1.76 – 2.50	1 - 1.75



แปลผลค่าเฉลี่ย

เนื้อหา/รูปแบบการฝึกอบรม	ระดับความพึงพอใจ (ร้อยละ)			
	มากที่สุด	มาก	น้อย	น้อยที่สุด
ค่าจริง	3	2	1	0
ค่าแปลผล	2.26 - 3	1.51 – 2.25	0.76 – 1.50	0 - 0.75



เกณฑ์การแปลผลค่าเฉลี่ย

เกณฑ์การแปลผล 5 ระดับ

- ค่าเฉลี่ย 4.21 - 5.00 - ระดับมากที่สุด
- ค่าเฉลี่ย 3.41 - 4.20 - ระดับมาก
- ค่าเฉลี่ย 2.61 - 3.40 - ระดับปานกลาง
- ค่าเฉลี่ย 1.81 - 2.60 - ระดับน้อย
- ค่าเฉลี่ย 1.00 - 1.80 - ระดับน้อยที่สุด

สูตรการคำนวณ

$(\text{Max-Min}) / \text{จำนวนชั้น}$

$$(5-1) / 5 = 4/5 = 0.80$$

เกณฑ์การแปลผล 3 ระดับ

- ค่าเฉลี่ยอยู่ระหว่าง 1- 2.33 - ระดับน้อย
- ค่าเฉลี่ย อยู่ระหว่าง 2.34-3.67 -ระดับปานกลาง
- ค่าเฉลี่ย อยู่ระหว่าง 3.68-5.00 -ระดับมาก

$$(5-1) / 3 = 4/3 = 1.33$$



ส่วนเบี่ยงเบนมาตรฐาน (SD)

- เป็นค่าวัดการกระจายที่นิยมใช้มากที่สุด เพราะใช้ค่าของข้อมูลทุกค่ามาคำนวณ
- ถ้าชุดข้อมูลมีการกระจายมาก ค่าสังเกตแต่ละค่าจะอยู่ห่างจากค่าเฉลี่ยมาก จึงทำให้ส่วนเบี่ยงเบนมาตรฐานมีค่ามาก
- ถ้าชุดข้อมูลมีการกระจายน้อย ค่าสังเกตแต่ละค่าเกาะกลุ่มอยู่ใกล้ ๆ ค่าเฉลี่ย จึงทำให้ส่วนเบี่ยงเบนมาตรฐานมีค่าน้อย



ตัวอย่างการวิเคราะห์สถานภาพองค์กร

ข้อคำถาม	Mean	SD	ระดับความคิดเห็น
หน่วยงานของท่านมีการจัดทำแผนอัตรากำลังซึ่งมีความเหมาะสมต่อพันธกิจหน่วยงาน	3.66	.99	ปานกลาง
หน่วยงานของท่านมีการวัดความพึงพอใจของบุคลากรอยู่เสมอ	3.46	1.09	ปานกลาง
หน่วยงานของท่านมีกระบวนการการจัดอบรมและพัฒนาบุคลากร	3.85	.95	มาก
ระบบการประเมินผลงานบุคลากรมีความเหมาะสมและเป็นธรรม	3.47	1.24	ปานกลาง
รวม	3.65	.80	ปานกลาง



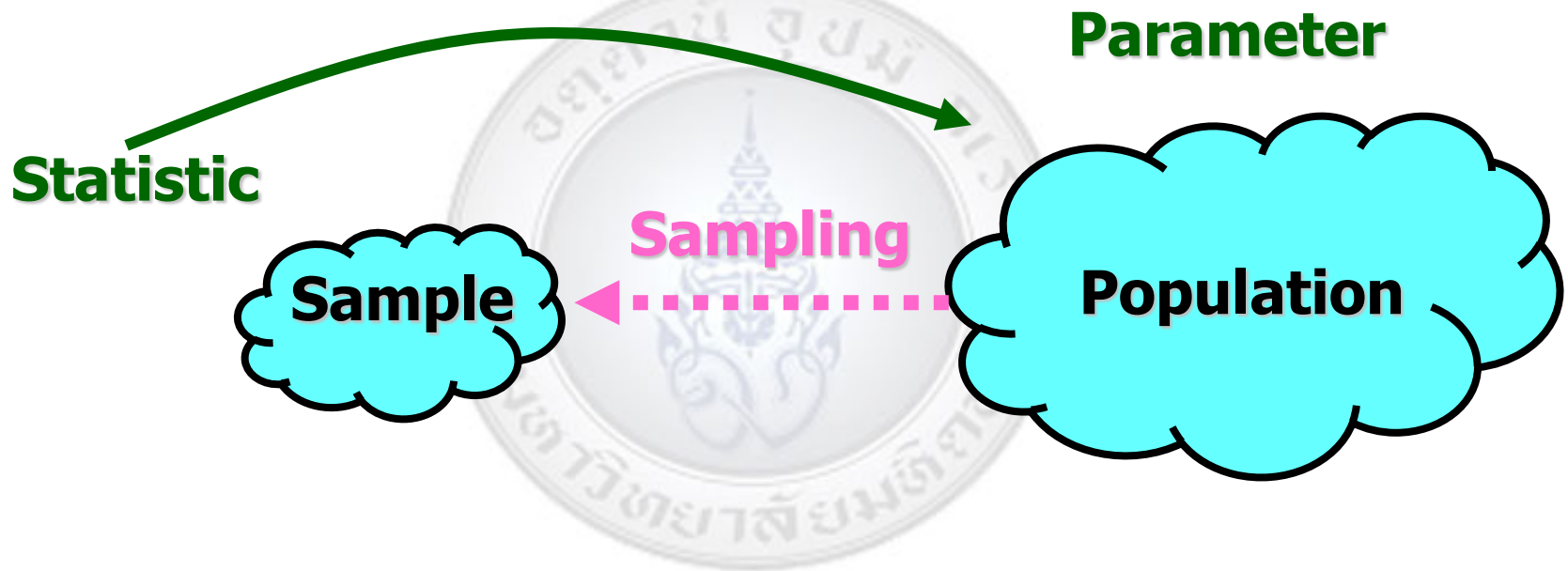
ตัวอย่างการวิเคราะห์ระดับความพึงพอใจ

ความพึงพอใจด้านเจ้าหน้าที่ บุคลากรผู้ให้บริการ	Mean	SD	ระดับความพึงพอใจ
ความสุภาพ กิริยามารยาทของเจ้าหน้าที่ผู้ให้บริการ	3.99	.68	มาก
ความเหมาะสมในการแต่งกาย บุคลิก ลักษณะท่าทางของเจ้าหน้าที่ผู้ให้บริการ	3.97	.66	มาก
ความเอาใจใส่ กระตือรือร้น และความพร้อมในการให้บริการของเจ้าหน้าที่	3.97	.67	มาก
เจ้าหน้าที่มีความรู้ ความสามารถในการให้บริการ เช่น การตอบคำถาม ชี้แจงข้อสงสัย ให้คำแนะนำ ช่วยแก้ปัญหาได้	3.95	.67	มาก
เจ้าหน้าที่ให้บริการต่อผู้รับบริการเหมือนกันทุกราย โดยไม่เลือกปฏิบัติ	3.97	1.43	มาก
ความซื่อสัตย์สุจริตในการปฏิบัติหน้าที่	3.96	.93	มาก
รวม	3.97	.54	มาก



สถิติที่ใช้ในการวิจัย

สถิติเชิงอนุมาน (Inference Statistics)





การวิเคราะห์ข้อมูลโดยใช้สถิติ t -Test

ลักษณะการทดสอบ : ใช้ทดสอบความแตกต่างระหว่างค่าเฉลี่ยของประชากร 2 กลุ่ม

ตัวอย่าง ต้องการทราบว่านักศึกษาชายและหญิงมีระดับความพึงพอใจต่อการให้บริการโดยเฉลี่ยแตกต่างกันหรือไม่



ลักษณะข้อมูล



มหาวิทยาลัยมหิดล
Mahidol University
"Wisdom of The Land"

มี 2 ตัวแปร คือ ตัวแปรตาม เป็นข้อมูลเชิงปริมาณ และตัวแปรอิสระ ที่เป็นข้อมูลเชิงคุณภาพ ที่แบ่งเป็น 2 กลุ่ม ที่เป็นอิสระจากกัน

เช่น ต้องการเปรียบเทียบระดับความพึงพอใจในงานระหว่างพนักงานเพศชาย และเพศหญิง ว่าแตกต่างกันหรือไม่ ที่ระดับนัยสำคัญทางสถิติ 0.05



ตัวอย่างการวิเคราะห์

ประเด็น	ความคาดหวัง		สภาพที่เป็นจริง		GAP	P-Value
	mean	S.D.	mean	S.D.		
ปัจจัยด้านกระบวนการ						
วัตถุประสงค์ของแต่ละหลักสูตรที่เรียนมีความชัดเจน	4.28	.72	4.22	.68	0.06	.022
เนื้อหาของหลักสูตรที่เรียนมีความสอดคล้องกับวัตถุประสงค์	4.28	.70	4.21	.71	0.07	.021
เนื้อหาในหลักสูตรที่เรียนไม่ยากเกินไป	4.20	.77	4.09	.74	0.11	.001*
การออกแบบเนื้อหา/บทเรียน น่าสนใจ เข้าใจง่าย และกระตุ้นให้ผู้เรียนเกิดความอยากเรียน	4.23	.77	4.07	.79	0.16	.000*
แต่ละหลักสูตรมีกิจกรรมเสริมเพื่อช่วยให้การเรียนรู้ง่ายขึ้น	4.20	.78	3.96	.83	0.24	.000*
มีการประเมินผลการเรียนเป็นระยะ ๆ	4.13	.78	3.91	.83	0.22	.000*
มีการปฏิสัมพันธ์กับผู้จัดหลักสูตรและผู้เรียนด้วยกัน	3.86	.96	3.33	1.01	0.53	.000*
สรุปด้านกระบวนการ	4.17	.66	3.97	.62	0.20	.000*



คะแนนความพึงพอใจ

เพศชาย(1)

50
51
56
54
52
54

เพศหญิง(2)

49
48
47
45
50
46



การลงข้อมูลใน Data Editor ใน SPSS

Sex = เพศ

1 คือ ชาย

2 คือ หญิง

Score = คะแนนความ
พึงพอใจในงาน

sex	score
1	50
1	51
1	56
1	54
1	52
1	54
2	49
2	48
2	47
2	45
2	50
2	46

File Edit View Data Transform Analyze Graphs Utilities Window Help

Reports
Descriptive Statistics
Tables
Compare Means
General Linear Model
Mixed Models
Correlate
Regression
Loglinear

Means...
One-Sample T Test...
Independent-Samples T Test...
Paired-Samples T Test...
One-Way ANOVA...

	sex	score
1	1	
2	1	
3	1	

Independent-Samples T Test

Test Variable(s):
score

Grouping Variable:
sex[? ?]

Define Groups...

Options.

OK
Paste
Reset
Cancel
Help

Define Groups

☒ Use specified values

Group 1: 1

Group 2: 2

☐ Cut point:

Continue
Cancel
Help

Group Statistics

SEX		N	Mean	Std. Deviation	Std. Error Mean
SCORE	1	6	52.83	2.229	.910
	2	6	47.50	1.871	.764

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
SCORE	Equal variances assumed	.385	.549	4.490	10	.001	5.33	1.188	2.687	7.980
	Equal variances not assumed			4.490	9.709	.001	5.33	1.188	2.676	7.991

$$H_0 : \mu_{male} = \mu_{female}$$

$$H_1 : \mu_{male} \neq \mu_{female}$$

เป็นการทดสอบแบบ 2 ทาง

P-Value = Sig.(2-tailed) = 0.001



Results

พนักงานเพศชาย มีความพอใจในงานโดยเฉลี่ย 52.83 คะแนน ในขณะที่พนักงานเพศหญิงมีความพอใจในงาน 47.50 คะแนน เมื่อทำการทดสอบสมมติฐานที่ระดับนัยสำคัญทางสถิติ 0.05 พบว่า พนักงานเพศชายมีความพอใจในงานโดยเฉลี่ยแตกต่างจากพนักงานเพศหญิง อย่างมีนัยสำคัญทางสถิติ



การวิเคราะห์ข้อมูลโดยใช้สถิติ F - Test (One Way ANOVA)

ลักษณะการทดสอบ : ใช้ทดสอบความแตกต่าง
ระหว่างค่าเฉลี่ยของประชากรหลายกลุ่ม
(มากกว่า 2 กลุ่ม)

ตัวอย่าง ต้องการทราบว่าช่วงอายุที่แตกต่างกัน
ส่งผลต่อความพึงพอใจในการฝึกอบรมต่างกัน
หรือไม่



คะแนนความพึงพอใจในงาน

ผู้บริหารระดับสูง(1) ผู้บริหาร
ระดับกลาง(2) ผู้บริหาร
ระดับต้น(3)

5
4
6
5

6
7
5
7

8
9
9
8



การลงข้อมูล

level = ระดับ

1 คือ สูง

2 คือ กลาง

3 คือ ต่ำ

Score = คะแนนความ
พอใจในงาน

	level	score
1	1	5
2	1	4
3	1	6
4	1	5
5	2	6
6	2	7
7	2	5
8	2	7
9	3	8
10	3	9
11	3	9
12	3	8

File Edit View Data Transform Analyze Graphs Utilities Window Help

8 :

	level	score
1	1	
2	1	
3	1	

Analyze

- Reports
- Descriptive Statistics
- Tables
- Compare Means**
 - Means...
 - One-Sample T Test...
 - Independent-Samples T Test...
 - Paired-Samples T Test...
 - One-Way ANOVA...**
- General Linear Model
- Mixed Models
- Correlate
- Regression

One-Way ANOVA

Dependent List: # score

Factor: # level

Contrasts... Post Hoc...

OK

One-Way ANOVA: Options

Statistics

- ☒ Descriptive
- ☐ Fixed and random effects
- ☐ Homogeneity of variance test
- ☐ Brown-Forsythe
- ☐ Welch

Means plot

- ☐ Means plot

Missing Values

- ☒ Exclude cases analysis by analysis

Continue Cancel Help

Descriptives

SCORE

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
1	4	5.00	.816	.408	3.70	6.30	4	6
2	4	6.25	.957	.479	4.73	7.77	5	7
3	4	8.50	.577	.289	7.58	9.42	8	9
Total	12	6.58	1.676	.484	5.52	7.65	4	9

ANOVA

SCORE

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	25.167	2	12.583	19.696	.001
Within Groups	5.750	9	.639		
Total	30.917	11			



Results

กรณี การวิเคราะห์ความแปรปรวน (ANOVA)
พบว่า มีนัยสำคัญทางสถิติ ต้องทำการ
เปรียบเทียบค่าเฉลี่ยรายคู่ เพื่อดูว่าค่าเฉลี่ย
คู่ใด แตกต่างหรือไม่แตกต่างกัน

One-Way ANOVA [X]

▶

Dependent List:
score

▶

Factor:
level

Contrasts...

Post Hoc...

Options...

OK

Paste

Reset

Cancel

Help

One-Way ANOVA: Post Hoc Multiple Comparisons [X]

Equal Variances Assumed

☒ LSD	☐ S-N-K	☐ Waller-Duncan
☐ Bonferroni	☐ Tukey	Type I/Type II Error Ratio: 100
☐ Sidak	☐ Tukey's-b	☐ Dunnett
☐ Scheffe	☐ Duncan	Control Category: Last
☐ R-E-G-W F	☐ Hochberg's GT2	Test
☐ R-E-G-W Q	☐ Gabriel	☒ 2-sided ☐ < Control ☐ > Control

Equal Variances Not Assumed

☐ Tamhane's T2	☐ Dunnett's T3	☐ Games-Howell	☐ Dunnett's C
---------------------------------------	---------------------------------------	---------------------------------------	--------------------------------------

Significance level: .05

Continue

Cancel

Help

Multiple Comparisons

Dependent Variable: SCORE

LSD

(I) LEVEL	(J) LEVEL	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
1	2	-1.25	.565	.054	-2.53	.03
	3	-3.50*	.565	.000	-4.78	-2.22
2	1	1.25	.565	.054	-.03	2.53
	3	-2.25*	.565	.003	-3.53	-.97
3	1	3.50*	.565	.000	2.22	4.78
	2	2.25*	.565	.003	.97	3.53

*. The mean difference is significant at the .05 level.





Results

ผู้บริหารระดับสูง มีความพอใจในงานโดยเฉลี่ย 5.00 คะแนน ในขณะที่ผู้บริหารระดับกลาง และระดับต้นมีความพอใจในงานโดยเฉลี่ย 6.25 และ 8.50 คะแนน ตามลำดับ เมื่อทำการทดสอบสมมติฐานที่ระดับนัยสำคัญทางสถิติ 0.05 พบว่า มีความแตกต่างกันระหว่างกลุ่ม และผลจากการทดสอบค่าเฉลี่ยรายคู่โดยวิธี LSD พบว่า ผู้บริหารระดับต้นจะพอใจในงานมากที่สุด ส่วนผู้บริหารระดับกลาง และสูงมีความพอใจในงานไม่แตกต่างกัน



การวิเคราะห์ข้อมูลโดยใช้สถิติ Chi-Square Test

ลักษณะการทดสอบ : ใช้ทดสอบความสัมพันธ์
ระหว่างประชากร 2 กลุ่ม ที่เป็นอิสระต่อกัน
ตัวอย่าง ต้องการทราบว่าตำแหน่งงานมี
ความสัมพันธ์กับรูปแบบการฝึกอบรมหรือไม่



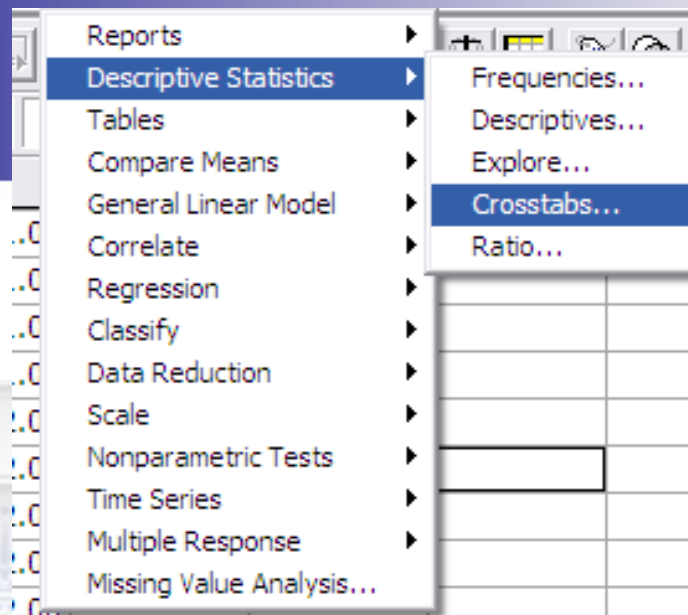
Chi-Square Test: χ^2

- สัมภาษณ์นักเรียน 100 คน จากอาจารย์ 3 ท่าน

ผลการ สัมภาษณ์	อาจารย์ท่านที่			รวม
	1	2	3	
ผ่าน	22	28	30	80
ไม่ผ่าน	8	2	10	20
ผลรวม	30	30	40	100



	กิจกรรมการ	ผล
19	1.00	1.00
20	1.00	1.00
21	1.00	1.00
22	1.00	1.00
23	1.00	2.00
24	1.00	2.00
25	1.00	2.00
26	1.00	2.00
27	1.00	2.00
28	1.00	2.00
29	1.00	2.00
30	1.00	2.00
31	2.00	1.00



Crosstabs: Statistics

☒ Chi-square

Nominal

☐ Contingency coefficient

☐ Phi and Cramér's V

☐ Lambda

☐ Uncertainty coefficient

ผล * กรรมการ Crosstabulation

Count

		กรรมการ			Total
		คนที่ 1	คนที่ 2	คนที่ 3	
ผล	ผ่าน	22	28	30	80
	ไม่ผ่าน	8	2	10	20
Total		30	30	40	100

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	4.792 ^a	2	.091
Likelihood Ratio	5.603	2	.061
Linear-by-Linear Association	.000	1	1.000
N of Valid Cases	100		

Asymp.Sig. (0.091) มีค่ามากกว่าระดับนัยสำคัญ(0.01) แสดงว่า
ไม่มีนัยสำคัญทางสถิติ สรุปได้ว่า ผลการสัมภาษณ์ของอาจารย์
ทั้ง 3 ท่าน ไม่มีความสัมพันธ์กัน



สรุปเทคนิคทางสถิติ

ตัวแปรอิสระ	ตัวแปรตาม	สถิติที่ใช้
ตัวแปรเชิงกลุ่ม 1 ตัวที่สนใจ (ที่มีค่าย่อย 2 ค่า) เช่น เพศ	ตัวแปรเชิงปริมาณ 1 ตัว เช่น คะแนนระดับความพึงพอใจ	t -Test
ตัวแปรเชิงกลุ่ม 1 ตัวที่สนใจ (ที่มีค่าย่อยมากกว่า 2 ค่า) เช่น สถานภาพสมรส	ตัวแปรเชิงปริมาณ 1 ตัว เช่น คะแนนระดับความพึงพอใจ	F -Test (One way ANOVA)
ตัวแปรเชิงปริมาณ	ตัวแปรเชิงปริมาณ	Regression หรือ Correlation



R2R วิเคราะห์เสร็จก็เตรียมผลงาน เพื่อการตีพิมพ์ในวารสารวิชาการ