

บทที่ 3

วิธีดำเนินการวิจัย

การวิจัยครั้งนี้มีวัตถุประสงค์หลักเพื่อศึกษาและเปรียบเทียบประสิทธิภาพของวิธีการประมาณค่าพารามิเตอร์ 4 วิธี ได้แก่ 1) Maximum Likelihood (ML) 2) Maximum Likelihood (ML) ในกรณีที่มีการเปลี่ยนรูปข้อมูลให้มีการแจกแจงเป็นโค้งปกติ 3) Robust Weighted Least Square – mean and Variance Adjusted χ^2 (WLSMV) และ 4) Bayesian สำหรับงานวิจัยที่ใช้เทคนิคการวิเคราะห์องค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis) โดยข้อมูลประชากรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็น โค้งปกติ (Non – normal Ordered Categorical Data) โดยใช้การวิเคราะห์ข้อมูลทุติยภูมิ (Secondary Analysis) ที่อ้างอิงจากงานวิจัยเชิงสำรวจ ซึ่งผลการวิจัยจะได้สารสนเทศเกี่ยวกับแนวการปฏิบัติในการวิเคราะห์ข้อมูลโมเดลองค์ประกอบเชิงยืนยัน ซึ่งมีรายละเอียดการวิจัยดังนี้

ฐานข้อมูลที่ใช้ในการวิจัย

ประชากร

ข้อมูลประชากรที่ใช้ในการศึกษาในครั้งนี้ เป็นข้อมูลประชากรเทียมที่สร้างขึ้น โดยอ้างอิงและประยุกต์จากงานวิจัยเชิงสำรวจ ที่มีลักษณะข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) ที่ใช้มาตรวัด Likert Type Scale 5 ระดับ จำนวน 10,000 ชุดข้อมูล และมีลักษณะการแจกแจงในแต่ละตัวแปร (Univariate Distribution) ไม่เป็น โค้งปกติ คือ เบ้ทางซ้าย (Negative Skewness) ทุกตัวแปร และไม่มีข้อมูลขาดหาย ซึ่งค่าสถิติพื้นฐานของข้อมูลประชากร การทดสอบการแจกแจงปกติของแต่ละตัวแปร (Univariate Normality Test) และการทดสอบการแจกแจงปกติพหุ (Multivariate Normality Test) โดยพิจารณาจากนัยสำคัญของค่ามาตรฐาน (Z – score และ ค่า p – value) ของค่าความเบ้และค่าความโด่ง และค่า Chi – square ของความเบ้และความโด่ง รายละเอียดดังตารางที่ 3 – 1

ตารางที่ 3 – 1 ค่าเฉลี่ย ค่าภาคเคลื่อนมาตรฐานของค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน การทดสอบการแจกแจงปกติของแต่ละตัวแปร การทดสอบการแจกแจงปกติของข้อมูลประชากร ($N = 10,000$)

Variables	$\bar{X}$	$SE_{\bar{X}}$	SD	Test of Univariate Normality						Skewness and Kurtosis	
				Skewness		Kurtosis		χ^2	p		
				Value	Z-score	Value	Z-score				
A1	4.084	.007	.710	-.532	-9.803	.000	.596	-4.917	.000	120.272	.000
A2	4.022	.008	.788	-.633	-10.803	.000	.601	-7.564	.000	173.923	.000
A3	4.012	.007	.747	-.502	-8.967	.000	.450	-6.253	.000	119.501	.000
A4	3.967	.008	.782	-.566	-9.171	.000	.540	-6.116	.000	121.513	.000
B1	3.823	.008	.770	-.429	-5.247	.000	.574	-4.010	.000	43.606	.000
B2	3.871	.008	.796	-.403	-6.187	.000	.301	-7.467	.000	94.037	.000
B3	3.946	.008	.781	-.423	-7.929	.000	.178	-7.959	.000	126.213	.000
B4	3.952	.007	.746	-.393	-7.315	.000	.222	-5.692	.000	85.912	.000
Test of Multivariate Normality				2.130	48.675	.000	106.653	58.378	.000	5777.197	.000

จากตารางที่ 3 – 1 การทดสอบการแจกแจงปกติของแต่ละตัวแปร (Univariate Normality Test) และการทดสอบการแจกแจงปกติพหุ (Multivariate Normality Test) ของข้อมูลประชากร พบว่า มีการแจกแจงที่แตกต่างไปจากโค้งปกติอย่างมีนัยสำคัญทางสถิติ ($p < .001$) จึงสามารถสรุปได้ว่าข้อมูลประชากรมีลักษณะการแจกแจงไม่เป็นโค้งปกติ (Non – normality Distribution)

กลุ่มตัวอย่าง

กลุ่มตัวอย่างได้จากสุ่มแบบใส่คืนจากประชากร ($N = 10,000$) ด้วยขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด (80, 160, 370 และ 623 หน่วยตัวอย่าง) โดยทำการศึกษาในแต่ละเงื่อนไขจำนวน 200 การทดลองซ้ำ (Replications)

การกำหนดโมเดลในการวิเคราะห์

โมเดลที่ใช้เป็นต้นแบบในการศึกษา ผู้วิจัยดำเนินการสร้าง โมเดลการวัด หรือโมเดลการวิเคราะห์ห่อองค์ประกอบเชิงยืนยันลำดับแรก (First Order Confirmatory Factor Analysis) ที่ประกอบด้วย 2 องค์ประกอบ องค์ประกอบละ 4 ตัวแปรสังเกตได้ รวมทั้งสิ้น 8 ตัวแปรสังเกตได้ โดยดำเนินการสกัดองค์ประกอบด้วยวิธีการวิเคราะห์ห่อองค์ประกอบเชิงสำรวจ (Exploratory Factor Analysis: EFA) ซึ่งมีรายละเอียดดังต่อไปนี้

ตารางที่ 3 – 2 การระบุจำนวนองค์ประกอบจากวิธีการวิเคราะห์ห่อองค์ประกอบเชิงสำรวจ

Variables	Component 1	Component 2
A1	.761	.245
A2	.838	.191
A3	.840	.204
A4	.808	.159
B1	.268	.760
B2	.173	.857
B3	.187	.832
B4	.160	.727

Eigen Values = 4.021 และ 1.477, Percent of Variance Criterion = 68.738%, Kaiser

Meyer – Olkin Measure of Sampling Adequacy = .857, Bartlett's Test ($p < .001$)

จากตารางที่ 3 – 2 พบว่า ตัวแปรสังเกตได้สามารถระบุจำนวนองค์ประกอบได้เป็น 2 องค์ประกอบ เพราะค่า Eigen Values ที่มากกว่า 1.000 มี 2 องค์ประกอบ ค่า Percent of Variance Criterion หรือค่าร้อยละของความแปรปรวนทั้งหมดที่อธิบายได้ (Total Variance Explained) มีค่าเท่ากับ 68.738% ซึ่งในงานวิจัยทางสังคมศาสตร์จะถือเกณฑ์การสกัดองค์ประกอบเมื่อร้อยละความแปรปรวนสะสมไม่ต่ำกว่า 60% ค่าดัชนี Kaiser Meyer Olkin (KMO) มีค่าเท่ากับ .857 เข้าใกล้ 1.000 ถือว่ามีความสัมพันธ์กันระหว่างตัวแปรมากพอ (Measure of Sampling Adequacy) ที่จะนำมาวิเคราะห์องค์ประกอบ ส่วนค่าสถิติ Bartlett's Test of Sphericity มีนัยสำคัญทางสถิติ ($p < .001$) แสดงว่าเมตริกซ์สหสัมพันธ์ของประชากรไม่เป็นเมตริกซ์เอกลักษณะ เมตริกซ์สหสัมพันธ์นั้นมีความเหมาะสมที่จะใช้วิเคราะห์องค์ประกอบต่อไป (Bollen, 1989)

การวิเคราะห์โมเดลองค์ประกอบเชิงยืนยันของโมเดลพารามิเตอร์

เนื่องจากเงื่อนไขข้อตกลงเบื้องต้นเกี่ยวกับลักษณะข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็น โค้งปกติ (Non – normality Ordered Categorical Data) ผู้วิจัยดำเนินการวิเคราะห์องค์ประกอบเชิงยืนยัน ด้วยวิธีการหาค่าพารามิเตอร์ใน โมเดลประชากรเพื่อใช้เป็นต้นแบบในการเปรียบเทียบค่าประมาณของพารามิเตอร์จากการวิเคราะห์ข้อมูลจากกลุ่มตัวอย่าง โดยแบ่งออกเป็น 2 โมเดล ตามฐานแนวคิดของวิธีการประมาณค่าดังนี้

โมเดลพารามิเตอร์ที่ 1 การหาค่าพารามิเตอร์ด้วยวิธี WLSMV เพื่อใช้เป็นต้นแบบในการเปรียบเทียบการประมาณค่าพารามิเตอร์จากกลุ่มตัวอย่างด้วยวิธีการ ML, ML_r และ WLSMV

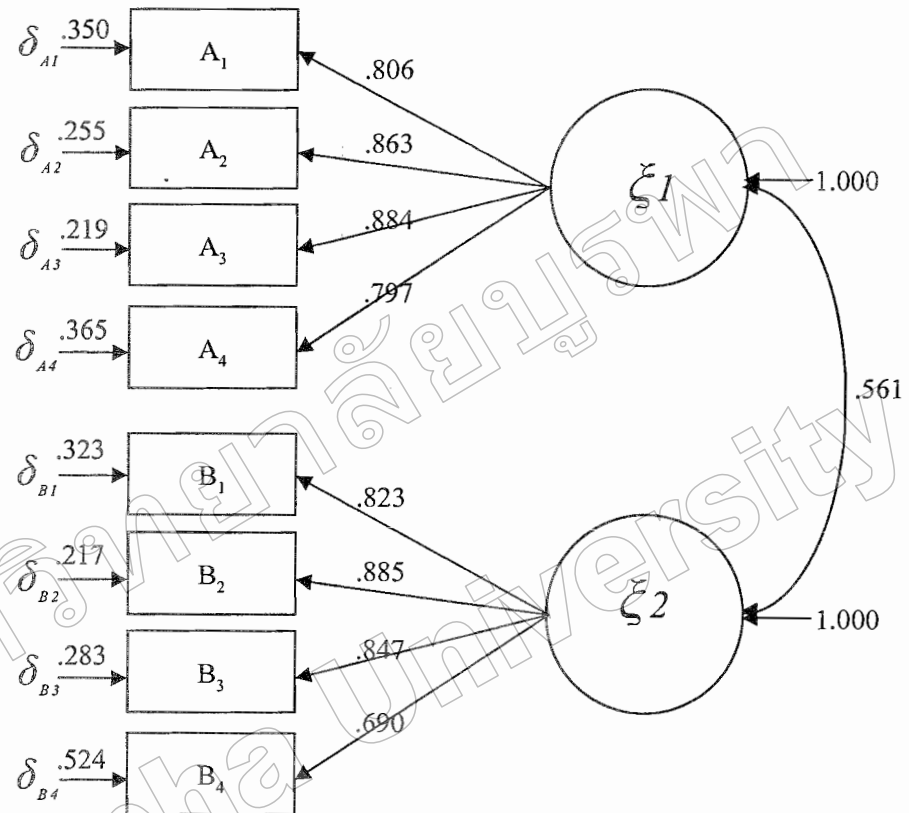
โมเดลพารามิเตอร์ที่ 2 การหาค่าพารามิเตอร์ด้วยวิธี Bayesian เพื่อใช้เป็นต้นแบบในการเปรียบเทียบการประมาณค่าพารามิเตอร์จากกลุ่มตัวอย่างด้วยวิธีการ Bayesian

ตารางที่ 3 – 3 ค่าพหุนัยองค์ประกอบ ค่าคลาดเคลื่อนมาตรฐาน ค่า Square Multiple Correlation และความคลาดเคลื่อนจากการวัด ในรูปแบบมาตรฐาน
ของโมเดลพารามิเตอร์ ($N = 10,000$)

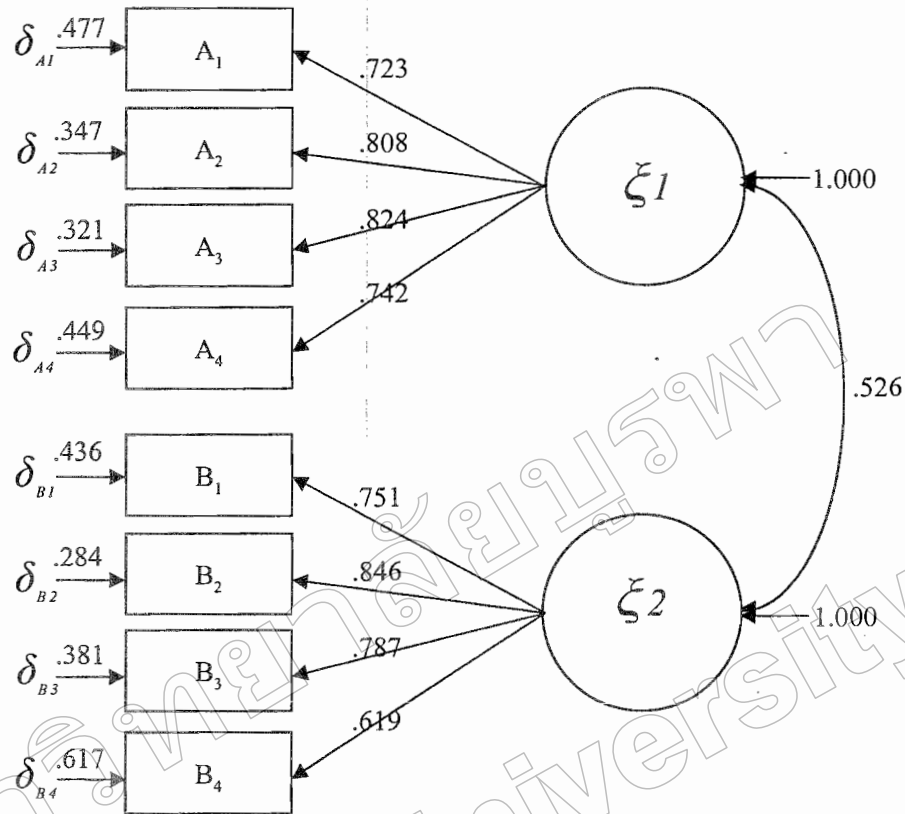
Parameters	Parameter Model 1 (WLSMV)				Parameter Model 2 (Bayesian)			
	Loading	Standard Error	Square Multiple Correlation	Measurement Error	Loading	Standard Error	Square Multiple Correlation	Measurement Error
λ_{A1}	.806*	.004	.650	.350	.723*	.006	.523	.477
λ_{A2}	.863*	.004	.745	.255	.808*	.004	.653	.347
λ_{A3}	.884*	.003	.781	.219	.824*	.005	.679	.321
λ_{A4}	.797*	.005	.635	.365	.742*	.006	.551	.449
λ_{B1}	.823*	.004	.677	.323	.751*	.005	.564	.436
λ_{B2}	.885*	.003	.783	.217	.846*	.004	.716	.284
λ_{B3}	.847*	.004	.717	.283	.787*	.006	.619	.381
λ_{B4}	.690*	.006	.476	.524	.619*	.007	.383	.617
Φ	.561*	.008			.526*	.008		

* $p < .001$

จากตารางที่ 3-3 สามารถแสดงรายละเอียดค่าพารามิเตอร์ในโมเดลความสัมพันธ์ระหว่างตัวแปรสังเกตได้และตัวแปรแฝงได้ดังภาพที่ 3-1 และ ภาพที่ 3-2



ภาพที่ 3-1 โมเดลพารามิเตอร์ 1 ค่าพารามิเตอร์จากการคำนวณด้วยวิธี WLSMV



ภาพที่ 3 – 2 โมเดลพารามิเตอร์ 2 ค่าพารามิเตอร์จากการคำนวณด้วยวิธี Bayesian

ขั้นตอนการวิจัย

1. ทดสอบการแจกแจงปกติของแต่ละตัวแปร (Univariate Normality Test) และการทดสอบการแจกแจงปกติพหุ (Multivariate Normality Test) ของข้อมูลประชากร ซึ่งพบว่ามีการแจกแจงไม่เป็น โค้งปกติ (Non – normality Distribution)
2. คำนวณค่าพารามิเตอร์ด้วยวิธีการ WLSMV โดยกำหนดให้เป็นโมเดลพารามิเตอร์ที่ 1 เพื่อใช้เปรียบเทียบการประมาณค่าพารามิเตอร์จากกลุ่มตัวอย่างด้วยวิธีการ ML, ML_{Tr} และ WLSMV และคำนวณค่าพารามิเตอร์ด้วยวิธีการ Bayesian โดยกำหนดให้เป็นโมเดลพารามิเตอร์ที่ 2 เพื่อใช้เปรียบเทียบการประมาณค่าพารามิเตอร์จากกลุ่มตัวอย่างด้วยวิธีการ Bayesian
3. กลุ่มข้อมูลจากประชากร ($N = 10,000$) ภายได้เงื่อนไขกลุ่มตัวอย่างที่ต่างกัน 4 ขนาด คือ 80, 160, 370 และ 623 หน่วยตัวอย่าง เงื่อนไขละ 200 การทดลองซ้ำ (Replications)
4. ดำเนินการประมาณค่าพารามิเตอร์ในแต่ละเงื่อนไข โดยการกำหนดโมเดล และการปรับโมเดล จะดำเนินการเหมือนกับโมเดลพารามิเตอร์ทุกประการ ซึ่งวิธีการประมาณค่าพารามิเตอร์ 4 วิธี ประกอบด้วย

4.1 ML เป็นวิธีการประมาณค่าพารามิเตอร์ภายใต้การละเมิดข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงแบบ โค้งปกติ และข้อมูลแบบต่อเนื่อง

4.2 ML_r เป็นวิธีการเปลี่ยนรูปข้อมูล (Data Transforming) ในแต่ละตัวแปร โดยวิธีการคำนวณ Normal Score ด้วยโปรแกรม LISREL Version 8.7 ซึ่งทำให้ข้อมูลมีการแจกแจงเป็น โค้งปกติและมีลักษณะเป็นข้อมูลแบบต่อเนื่องตามข้อตกลงเบื้องต้นเสียก่อนแล้วดำเนินการประมาณค่าพารามิเตอร์ด้วยวิธีการ ML

4.3 WLSMV เป็นวิธีการประมาณค่าพารามิเตอร์ที่มีความแข็งแกร่งต่อการละเมิดข้อตกลงเบื้องต้นเกี่ยวกับลักษณะข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็น โค้งปกติ และกลุ่มตัวอย่างขนาดเล็ก

4.4 Bayesian เป็นวิธีการประมาณค่าพารามิเตอร์ที่มีความแข็งแกร่งต่อการละเมิดข้อตกลงเบื้องต้นเกี่ยวกับลักษณะข้อมูลจำแนกกลุ่ม (Categorical Data) สามารถให้ค่าประมาณได้ในทุกรูปแบบการแจกแจง โดยค่าพารามิเตอร์ที่ได้และสถิติทดสอบมีความแข็งแกร่งต่อการแจกแจงที่ไม่เป็น โค้งปกติ และมีความแข็งแกร่งต่อกลุ่มตัวอย่างขนาดเล็ก สามารถประมาณค่าได้ดีโดยไม่มีข้อจำกัดเรื่องกลุ่มตัวอย่างขนาดใหญ่ ระยะเวลาในการคำนวณ โดยคอมพิวเตอร์มีความรวดเร็วกว่าวิธีการอื่น ๆ (Browne & Draper, 2006, p. 505 cited in Muthén & Asparouhov, 2011, p. 6) อีกทั้งยังสามารถวิเคราะห์ข้อมูลสำหรับ โมเดลแบบใหม่ ๆ ได้ดี เช่น โมเดลที่มีจำนวนพารามิเตอร์มาก ๆ

5. เปรียบเทียบประสิทธิภาพของวิธีการประมาณพารามิเตอร์ตามเงื่อนไขที่แตกต่างกัน รวมทั้งสิ้น จำนวน $4 \times 4 = 16$ เงื่อนไข (ขนาดกลุ่มตัวอย่าง $\times$ วิธีการประมาณค่า) โดยนำผลลัพธ์ที่ได้จาก 200 ชุดการทดลองซ้ำ (Replications) ในแต่ละเงื่อนไขมาหาค่าเฉลี่ย ดังนี้

5.1 ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าประมาณพารามิเตอร์ (Relative Parameter Estimates Bias)

5.2 ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias)

5.3 ค่าความเชื่อมั่น (Reliability) ของโมเดลการวัด ได้แก่ ค่า ρ_c และค่า ρ_v

5.4 ร้อยละของ โมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit) โดยพิจารณาจากดัชนีวัดความสอดคล้องกลมกลืนของโมเดล (Model Fit Indices)

การวิเคราะห์ข้อมูลและสถิติที่ใช้ในการวิเคราะห์ข้อมูล

ผู้วิจัยดำเนินการวิเคราะห์ข้อมูลโดยใช้โปรแกรมสำเร็จรูป SPSS Version 18.0 และนำเสนอผลการวิเคราะห์ข้อมูล โดยแบ่งเป็น 6 ตอน ดังนี้

ตอนที่ 1 วิเคราะห์ค่าสถิติพื้นฐานของค่าประมาณพารามิเตอร์ (Parameter Estimates) ที่คำนวณโดยวิธีการประมาณค่าพารามิเตอร์แบบ ML, ML_{tr}, WLSMV และ Bayesian ในขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด (80, 160, 370 และ 623 หน่วยตัวอย่าง) ด้วยการนำเสนอค่าเฉลี่ย ($\bar{X}$) ส่วนเบี่ยงเบนมาตรฐาน (SD) ค่าคลาดเคลื่อนมาตรฐานของค่าเฉลี่ย ($SE_{\bar{X}}$) และช่วงความเชื่อมั่น 95% แล้วพิจารณาเปรียบเทียบว่าค่าทำนายของค่าประมาณพารามิเตอร์ (Parameter Estimates) ในช่วงความเชื่อมั่น 95% ของแต่ละเงื่อนไข มีค่าครอบคลุมค่าจริงพารามิเตอร์ (True Parameter) หรือไม่

ตอนที่ 2 การวิเคราะห์ความแตกต่างระหว่างค่าจริงของพารามิเตอร์ (True Parameters) และค่าประมาณของพารามิเตอร์ (Parameter Estimates) จากวิธีการประมาณค่าพารามิเตอร์แบบ ML, ML_{tr}, WLSMV และ Bayesian และขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด (80, 160, 370 และ 623 หน่วยตัวอย่าง) โดยนำค่าประมาณพารามิเตอร์ (Parameter Estimates) ลบด้วยค่าจริงของพารามิเตอร์ (True Parameters) ในแต่ละเงื่อนไข แล้วทดสอบความแตกต่างด้วยสถิติ One Sample t - test เทียบกับค่าอ้างอิงเท่ากับศูนย์

ตอนที่ 3 การวิเคราะห์ค่าร้อยละความเอนเอียงสัมพัทธ์ (Relative Bias) ของค่าประมาณพารามิเตอร์จากวิธีการประมาณค่าพารามิเตอร์แบบ ML, ML_{tr}, WLSMV และ Bayesian และขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด (80, 160, 370 และ 623 หน่วยตัวอย่าง) โดยคำนวณได้จากสูตร ดังนี้ (Bandalos, 2006, p.401; Flora & Curran, 2004, p. 466; Trierweiler, 2009, p. 75)

$$RB(\hat{\theta}) = \left| \frac{\hat{\theta}_{ij} - \theta_i}{\theta_i} \right| \times 100\%$$

$\hat{\theta}_{ij}$ คือ ค่าประมาณจากกลุ่มตัวอย่างที่ j และค่าพารามิเตอร์ที่ i

θ_i คือ ค่าพารามิเตอร์จากประชากร

เกณฑ์สำหรับค่าความเอนเอียงสัมพัทธ์ (Relative Bias) ของค่าประมาณพารามิเตอร์ (Curran, West, & Finch, 1996; Bandalos, 2006, p.402; Flora & Curran, 2004, p. 473; Kaplan, 1989 cited in Trierweiler, 2009, p. 75) คือ

ไม่เกิน 5.00%	หมายถึง มีค่าความเอนเอียงในระดับต่ำ (Trivial Bias)
5.01% – 10.00%	หมายถึง มีค่าความเอนเอียงในระดับปานกลาง (Moderate Bias)
มากกว่า 10.00%	หมายถึง มีค่าความเอนเอียงในระดับสูง (Substantial Bias)

ตอนที่ 4 การวิเคราะห์ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) ของค่าประมาณพารามิเตอร์ จากวิธีการประมาณค่าพารามิเตอร์ แบบ ML, ML_r, WLSMV และ Bayesian และขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด (80, 160, 370 และ 623 หน่วยตัวอย่าง) โดยคำนวณได้จากสูตร ดังนี้ (Bandalos, 2006, p.403; Flora & Curran, 2004, p. 466; Trierweiler, 2009, p. 75)

$$RB(SE(\theta_i)) = \frac{SE(\theta_i)_j - SE(\theta_i)}{SE(\theta_i)} \times 100\%$$

$SE(\theta_i)$ คือ ค่าประมาณของความคลาดเคลื่อนมาตรฐานของพารามิเตอร์จากประชากร

$SE(\theta_i)_j$ คือ ค่าประมาณของความคลาดเคลื่อนมาตรฐานจากกลุ่มตัวอย่าง สำหรับ

จำนวนชุดข้อมูล (Replication) ที่ j

เกณฑ์สำหรับค่าความเอนเอียงสัมพัทธ์ (Relative Bias) ของความคลาดเคลื่อนมาตรฐาน (Curran, West, & Finch, 1996; Bandalos, 2006, p.402; Flora & Curran, 2004, p. 473; Kaplan, 1989 cited in Trierweiler, 2009, p. 75) คือ

ไม่เกิน 5.00%	หมายถึง มีค่าความเอนเอียงในระดับต่ำ (Trivial Bias)
5.01% – 10.00%	หมายถึง มีค่าความเอนเอียงในระดับปานกลาง (Moderate Bias)
มากกว่า 10.00%	หมายถึง มีค่าความเอนเอียงในระดับสูง (Substantial Bias)

ตอนที่ 5 การเปรียบเทียบค่าความเชื่อมั่น (Reliability) ของโมเดลการวัด จากวิธีการประมาณค่าพารามิเตอร์แบบ ML, ML_r, WLSMV และ Bayesian และขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด (80, 160, 370 และ 623 หน่วยตัวอย่าง) ได้แก่

1. ความเชื่อมั่นของตัวแปรแฝง (Construct Reliability: ρ_c) แสดงถึงความตรงในการรวมตัว (Convergent Validity) ซึ่งหมายถึง สัดส่วนความแปรปรวนร่วมกันของตัวแปรสังเกตได้ทั้งหมดในตัวแปรแฝงเดียวกัน การคำนวณควรแยกแต่ละตัวแปรแฝง และค่าที่ได้ควรมากกว่า .60 โดยคำนวณ จากสูตร

$$\rho_c = \frac{(\sum \lambda_i)^2}{(\sum \lambda_i)^2 + \sum \delta_i}$$

เมื่อ λ_i คือ ค่ามาตรฐานของน้ำหนักองค์ประกอบ

δ_i คือ ค่ามาตรฐานของความแปรปรวนของความคลาดเคลื่อนจากการวัด

2. ค่าเฉลี่ยของความแปรปรวนที่ถูกสกัดได้ (Average Variance Extracted: ρ_v)

เป็นค่าเฉลี่ยของความแปรปรวนของตัวแปรแฝงที่อธิบายได้ด้วยตัวแปรสังเกตได้ ซึ่งมีค่าเทียบเท่ากับค่าไอเก้น (Eigen Value) ในการวิเคราะห์องค์ประกอบเชิงสำรวจ การคำนวณควรแยกแต่ละตัวแปรแฝง และค่าที่ได้ควรมากกว่า .50 ซึ่งคำนวณจากสูตร

$$\rho_v = \frac{(\sum \lambda_i)^2}{\sum \lambda_i^2 + \sum \delta_i}$$

เมื่อ λ_i คือ ค่ามาตรฐานของน้ำหนักองค์ประกอบ

δ_i คือ ค่ามาตรฐานความแปรปรวนของความคลาดเคลื่อนจากการวัด

ตอนที่ 6 การเปรียบเทียบร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit) โดยพิจารณาจากดัชนีวัดความสอดคล้องกลมกลืนของโมเดล (Model Fit Indices)

ค่าดัชนีวัดความสอดคล้องกลมกลืนของโมเดล ที่ได้จากวิธีการประมาณค่าพารามิเตอร์แบบ ML, ML_u และ WLSMV คือ ค่าดัชนี CFI (Comparative Fit Index) และค่าดัชนี TLI (Tucker – Lewis Index) หรือค่า NNFI (Non – normed Fit Index) ส่วนดัชนีวัดความสอดคล้องกลมกลืนของโมเดลจากวิธีการ Bayesian คือ ค่า Posterior Predictive p – value โดยยึดเกณฑ์ดังนี้

(Schemelleh, Moosbrugger & Müller, 2003, pp. 23 – 27)

CFI > .95 หมายถึง โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์

TLI > .95 หมายถึง โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์

Posterior Predictive p – value > .05 หมายถึง โมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์