

บทที่ 2

เอกสารและงานวิจัยที่เกี่ยวข้อง

ในการวิจัยครั้งนี้ผู้วิจัยได้ศึกษาเอกสารและงานวิจัยที่เกี่ยวข้อง และได้นำเสนอตามหัวข้อดังต่อไปนี้

1. โมเดลสมการเชิงโครงสร้าง (Structural Equation Modeling)
 - 1.1 ความหมายและความสำคัญของโมเดลสมการเชิงโครงสร้าง
 - 1.2 การตรวจสอบข้อมูลก่อนการวิเคราะห์โมเดลสมการเชิงโครงสร้าง
 - 1.3 ขั้นตอนการวิเคราะห์ข้อมูลด้วยโมเดลสมการเชิงโครงสร้าง
 - 1.4 ประเภทของพารามิเตอร์ในโมเดลสมการเชิงโครงสร้าง
2. โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis Model)
 - 2.1 ความหมายของโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน
 - 2.2 ขั้นตอนการวิเคราะห์องค์ประกอบเชิงยืนยัน
 - 2.3 การกำหนดประเภทของพารามิเตอร์และเมตริกซ์ที่ใช้ในการวิเคราะห์องค์ประกอบเชิงยืนยัน
 - 2.4 ขนาดกลุ่มตัวอย่างที่ใช้ในการวิเคราะห์องค์ประกอบเชิงยืนยัน และโมเดลสมการเชิงโครงสร้าง
3. วิธีการประมาณค่าพารามิเตอร์ (Parameter Estimation Method) ที่ใช้ในการวิจัย
4. ประสิทธิภาพของการประมาณค่าพารามิเตอร์ (The Efficiency of Parameter Estimates)
 - 4.1 ค่าความเอนเอียงของค่าประมาณพารามิเตอร์ (Parameter Estimates Bias)
 - 4.2 ค่าความเอนเอียงของค่าคลาดเคลื่อนมาตรฐาน (Standard Error Bias)
5. การประเมินความสอดคล้องของโมเดล (Assessment of Model Fit)
 - 5.1 การตรวจสอบโมเดลการวัด (Measurement Model)
 - 5.2 การตรวจสอบโมเดลโครงสร้าง (Structural Model)
 - 5.3 ดัชนีที่ใช้ในการตรวจสอบความสอดคล้องกลมกลืนระหว่างโมเดลตามสมมติฐานกับข้อมูลเชิงประจักษ์ (Fit Indices)
6. ผลกระทบจากการละเมิดข้อตกลงเบื้องต้นในการวิเคราะห์โมเดลสมการเชิงโครงสร้างและโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน
7. งานวิจัยที่เกี่ยวข้อง

โมเดลสมการเชิงโครงสร้าง (Structural Equation Modeling)

ความหมายและความสำคัญของโมเดลสมการเชิงโครงสร้าง

โมเดลสมการเชิงโครงสร้าง หรือ โมเดล SEM เป็นวิธีการทางสถิติอธิบาย ทดสอบและประมาณค่าความสัมพันธ์ของชุดข้อมูล วิเคราะห์โครงสร้างตามทฤษฎีที่รองรับเกี่ยวกับปรากฏการณ์หนึ่ง ๆ ที่กล่าวถึงความสัมพันธ์เชิงสาเหตุ (Causal) จากตัวแปรสังเกตได้หลายตัวแปรซึ่งมีลักษณะสำคัญ 2 ประการ คือ 1) เป็นการดำเนินการภายใต้การศึกษาข้อมูลที่ต้องอิงจากทฤษฎีเกี่ยวกับสมการเชิงโครงสร้าง และ 2) ความสัมพันธ์เชิงโครงสร้างสามารถแสดงด้วยแผนภาพโมเดลที่ชัดเจนตามกรอบแนวคิดตามทฤษฎีของการศึกษา (Byrne, 2012, p. 3; Schumacker & Lomax, 2010, p. 2) ซึ่งการสามารถทดสอบนัยสำคัญทางสถิติของโมเดลสมมติฐานสามารถทดสอบไปพร้อม ๆ กับการวิเคราะห์ข้อมูล คือ ถ้าดัชนีความสอดคล้องกลมกลืนของโมเดล (Goodness – of – fit) มีความเหมาะสม แสดงว่าโมเดลสมมติฐานตามทฤษฎีที่แสดงความสัมพันธ์ระหว่างตัวแปรมีความสอดคล้องกับข้อมูลเชิงประจักษ์ดี ถ้าไม่สอดคล้องหมายความว่า ปฏิเสธโมเดลสมมติฐาน

ลักษณะทั่วไปของโมเดลสมการเชิงโครงสร้างที่แตกต่างไปจากเทคนิคการวิเคราะห์ข้อมูลพหุตัวแปร (Multivariate Procedure) อื่น ๆ (Byrne, 2012, p. 3) คือ 1) ใช้วิเคราะห์เพื่อยืนยันโมเดลมากกว่าใช้วิเคราะห์เพื่อสำรวจหรือระบุโมเดล เหมาะสำหรับการทดสอบทฤษฎีมากกว่าการสร้างทฤษฎี 2) การวิเคราะห์พหุตัวแปรทั่วไป ไม่สามารถกำหนดหรือคำนวณความคลาดเคลื่อนจากการวัด (Measurement Error) ซึ่งส่วนมากมักไม่มีการอธิบายถึงความคลาดเคลื่อนดังกล่าว ซึ่งข้อผิดพลาดนี้อาจนำไปสู่ความไม่เหมาะสมที่เป็นปัญหาที่ร้ายแรงขึ้น โดยเฉพาะเมื่อขนาดของความคลาดเคลื่อนจากการวัดมีค่ามากขึ้น แต่การวิเคราะห์โมเดลสมการเชิงโครงสร้างสามารถคำนวณความแปรปรวนของความคลาดเคลื่อนค่าพารามิเตอร์ได้อย่างชัดเจน 3) การวิเคราะห์พหุตัวแปรทั่วไปใช้ฐานในการวัดเฉพาะตัวแปรที่สังเกตได้เท่านั้น แต่การวิเคราะห์โมเดลสมการเชิงโครงสร้างสามารถวัดความสัมพันธ์ของตัวแปรสังเกตได้และตัวแปรที่สังเกตไม่ได้ หรือตัวแปรแฝง (Latent Variables) ได้ในเวลาเดียวกัน และ 4) สามารถประมาณค่าโมเดลวิเคราะห์ความสัมพันธ์ตัวแปรพหุได้ทั้งแบบจุดและแบบช่วง รวมทั้งสามารถวิเคราะห์ความสัมพันธ์เชิงสาเหตุของตัวแปรได้ทั้งทางตรงและทางอ้อม (Schumacker & Lomax, 2010, pp. 2 – 3) ซึ่งจากลักษณะเด่นดังกล่าวทำให้การวิเคราะห์โมเดลสมการเชิงโครงสร้างเป็นเทคนิคที่ได้รับความนิยมอย่างแพร่หลายในการวิจัยทางสังคมศาสตร์ พฤติกรรมศาสตร์ การศึกษา จิตวิทยา ธุรกิจ และสาขาต่าง ๆ (Byrne, 2012, p. 4; Brown, 2006, p. 12)

ในการวิจัยด้านพฤติกรรมศาสตร์ นักวิจัยส่วนใหญ่มักสนใจศึกษาเกี่ยวกับตัวแปรที่ไม่สามารถวัดได้โดยตรงแต่มีโครงสร้างตามทฤษฎีแสดงผลออกมาในรูปของพฤติกรรมที่สามารถสังเกตได้ ซึ่งในการวิเคราะห์โมเดลสมการเชิงโครงสร้างจะเรียกว่า ตัวแปรแฝง (Latent Variables) เช่น ตัวแปรแฝงในงานวิจัยทางจิตวิทยา ได้แก่ ความเครียด แรงจูงใจ เป็นต้น งานวิจัยทางสังคมวิทยา ได้แก่ ลักษณะความเป็นมืออาชีพ (Professionalism) ลักษณะผิดปกติกทางสังคม (Anomie) เป็นต้น งานวิจัยทางการศึกษา ได้แก่ ความสามารถในการสอน ความคาดหวังของครู เป็นต้น งานวิจัยทางเศรษฐศาสตร์ ได้แก่ ระบบทุนนิยม (Capitalism) และระดับชั้นทางสังคม (Social Class) เป็นต้น ตัวแปรดังกล่าวล้วนแต่เป็นตัวแปรที่ไม่สามารถสังเกตได้โดยตรง ดังนั้นนักวิจัยจะศึกษาตัวแปรแฝงโดยการวัดตัวแปรพฤติกรรมที่สังเกตได้แทน เรียกว่า Observed Variables หรือ Manifest Variables หรือ Indicators ซึ่งจะอธิบายคุณลักษณะของตัวแปรแฝงดังกล่าว (Byrne, 2012, p. 4; Brown, 2006, pp. 2 – 3)

จากการศึกษาการวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรหลายตัว ไม่ว่าจะเป็นการวิเคราะห์สหสัมพันธ์พหุคูณ (Multiple Correlations) การวิเคราะห์ตัวแปรร่วม (Commonality Analysis) หรือการวิเคราะห์สหสัมพันธ์คาโนนิกอล (Canonical Correlation Analysis) ล้วนแต่ชี้ถึงความสัมพันธ์แบบธรรมดาระหว่างตัวแปรหรือกลุ่มตัวแปร ไม่ได้ยืนยันหรือสนับสนุนถึงความสัมพันธ์ในรูปที่เป็นสาเหตุและผล ซึ่งการยืนยันหรือสนับสนุนในรูปที่เป็นสาเหตุและผลก็คือการยืนยันหรือสนับสนุนว่าตัวแปรอิสระตัวใดเป็นสาเหตุให้เกิดความแปรปรวนหรือความแตกต่างในตัวแปรตาม และสาเหตุดังกล่าวเป็นสาเหตุที่เกิดจากตัวแปรอิสระนั้น ๆ โดยตรง หรือสาเหตุโดยทางอ้อม กล่าวคือ ไปร่วมกับตัวแปรอื่นในการทำให้เกิดความแปรปรวนในตัวแปรตามหรือเป็นไปทั้งสองทาง ซึ่งการวิเคราะห์โมเดลสมการเชิงโครงสร้างจึงเป็นการวิเคราะห์ข้อมูลที่มีโมเดลเชิงสาเหตุ (Causal Model)

การวิเคราะห์โมเดลเชิงสาเหตุเป็นการประยุกต์เพื่อการศึกษาความสัมพันธ์เชิงโครงสร้าง (Structural Relation) ระหว่างตัวแปรแฝงโดยการประยุกต์หลักการของ Path Analysis และ Confirmatory Factor Analysis: CFA (Schumacker & Lomax, 2010, pp. 4 – 5)

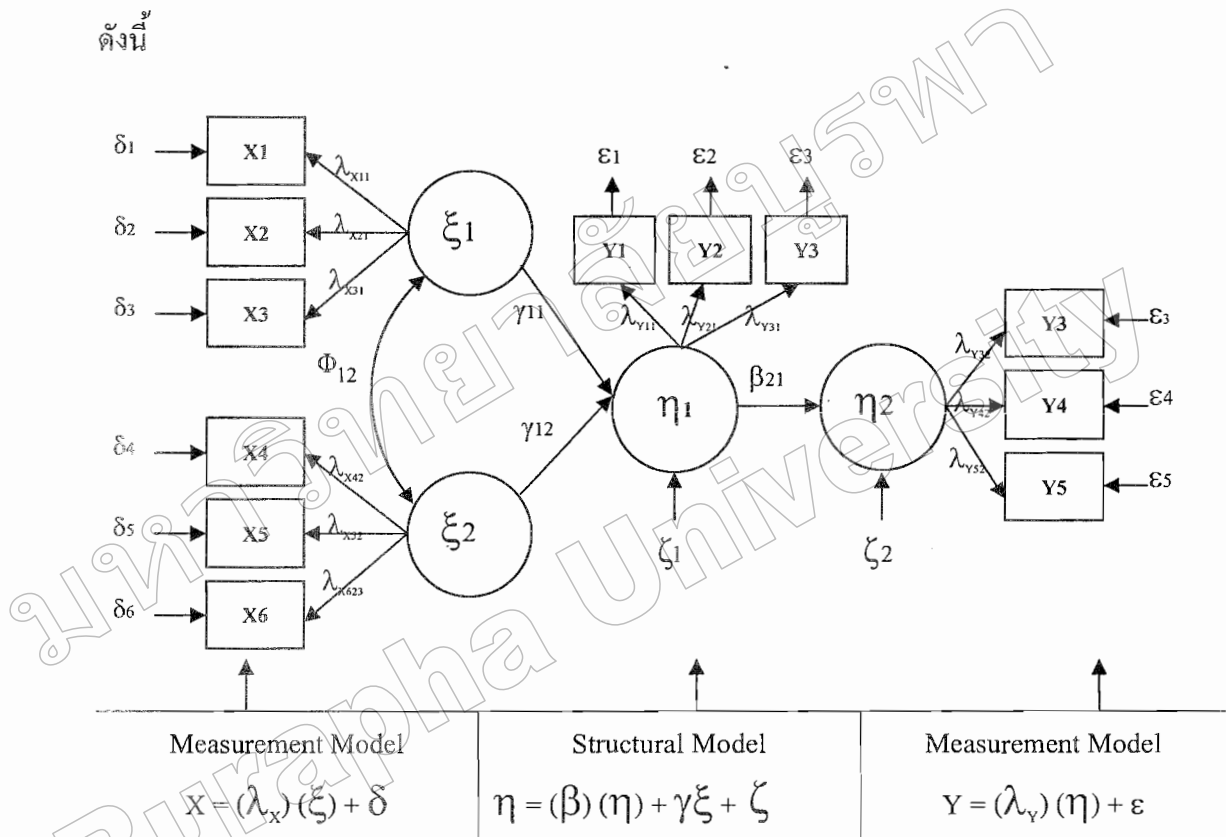
ตัวแปรที่ใช้ในโมเดลสมการเชิงโครงสร้าง (Byrne, 2012, p. 5)

ตัวแปรที่ใช้ในโมเดลสมการเชิงโครงสร้าง แบ่งเป็น 2 ประเภท คือ

1. ตัวแปรภายนอก (Exogenous Variables) หมายถึง ตัวแปรที่นักวิจัยไม่สนใจศึกษาสาเหตุของตัวแปรเหล่านี้ ตัวแปรสาเหตุของตัวแปรภายนอกจึงไม่ปรากฏในโมเดล
2. ตัวแปรภายใน (Endogenous Variables) หมายถึง ตัวแปรที่นักวิจัยสนใจศึกษาว่าได้รับอิทธิพลจากตัวแปรใด สาเหตุของตัวแปรภายในจะแสดงไว้ในโมเดลอย่างชัดเจน

ลักษณะทั่วไปของโมเดลสมการเชิงโครงสร้าง

ลักษณะของโมเดลสมการเชิงโครงสร้างประกอบด้วย 2 โมเดล คือ โมเดลการวัด (Measurement Model) เป็นโมเดลอธิบายความสัมพันธ์ระหว่างตัวแปรที่สังเกตได้กับตัวแปรแฝง และโมเดลสมการโครงสร้าง (Structural Equation Model) เป็นโมเดลอธิบายความสัมพันธ์ระหว่างตัวแปรแฝงด้วยกัน ซึ่งความสัมพันธ์ของตัวแปรในโมเดลทั้งหมด สามารถแสดงได้ตามภาพที่ 2 – 1 ดังนี้



ภาพที่ 2 – 1 โมเดลสมการเชิงโครงสร้าง

ตารางที่ 2 – 1 สัญลักษณ์ที่ใช้ในโมเดลสมการเชิงโครงสร้าง (Lee, 2007, pp. 6 – 7)

สัญลักษณ์กรีก	คำอ่าน	ภาษาอังกฤษ	ความหมาย
X	Eks	X	เวกเตอร์ตัวแปรภายนอกที่สังเกตได้
Y	Wi	Y	เวกเตอร์ตัวแปรภายในที่สังเกตได้
ξ	Ksi	K	เวกเตอร์ตัวแปรแฝงภายนอก
η	Eta	E	เวกเตอร์ตัวแปรแฝงภายใน
δ	Delta	d	เวกเตอร์ความคลาดเคลื่อน d ในการวัดตัวแปร x
ε	Epsilon	e	เวกเตอร์ความคลาดเคลื่อน e ในการวัดตัวแปร y
ζ	Zeta	z	เวกเตอร์ความคลาดเคลื่อน z ในการวัดตัวแปร η
λ_x	Lamda – X	LX	เมตริกซ์สัมประสิทธิ์การถดถอยของ X บน ξ
λ_y	Lamda – Y	LY	เมตริกซ์สัมประสิทธิ์การถดถอยของ Y บน η
γ	Gamma	GA	เมตริกซ์อิทธิพลเชิงสาเหตุจากตัวแปร ξ ต่อ η
β	Beta	BE	เมตริกซ์อิทธิพลเชิงสาเหตุระหว่างตัวแปร η
Φ	Phi	PH	เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมระหว่างตัวแปรภายนอกแฝง
Ψ	Psi	PS	เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมระหว่างความคลาดเคลื่อน z
Θ_δ	Theta – Delta	TD	เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมระหว่างความคลาดเคลื่อน d
Θ_ε	Theta – Epsilon	TE	เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมระหว่างความคลาดเคลื่อน e

โมเดลสมการเชิงโครงสร้าง ประกอบด้วยส่วนสำคัญ 2 ส่วน (Schumacker & Lomax, 2010, p. 184; Brown, 2006, p. 51) คือ

1. โมเดลการวัด (Measurement Model) เป็นโมเดลที่ระบุความสัมพันธ์เชิงเส้นระหว่างตัวแปรแฝงกับตัวแปรสังเกตได้ มี 2 ชนิด คือ โมเดลการวัดสำหรับตัวแปรแฝงภายนอก และโมเดลการวัดสำหรับตัวแปรแฝงภายใน หรือเป็นส่วนของการวิเคราะห์องค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis) จากภาพที่ 2 – 1 โมเดลการวัดสำหรับตัวแปรแฝงภายนอก

คือ โมเดลองค์ประกอบของ ζ_1 และ ζ_2 ส่วนโมเดลการวัดสำหรับตัวแปรแฝงภายใน คือ โมเดลองค์ประกอบของ η_1 และ η_2

2. โมเดลโครงสร้าง (Structural Model) เป็นโมเดลที่ระบุความสัมพันธ์ระหว่างตัวแปรแฝงกับตัวแปรแฝง

การตรวจสอบข้อมูล (Data Screening) ก่อนการวิเคราะห์โมเดลสมการเชิงโครงสร้าง

ในการดำเนินการวิจัยโดยใช้เทคนิคการวิเคราะห์โมเดลสมการเชิงโครงสร้างนั้น เป็นการศึกษาเกี่ยวกับความสัมพันธ์ของตัวแปร ซึ่งนักวิจัยควรที่รู้จักธรรมชาติของข้อมูลสำหรับการวิจัย และมีความเข้าใจลักษณะความสัมพันธ์ระหว่างตัวแปรเป็นอย่างดี จึงจะสามารถวิเคราะห์ข้อมูลได้อย่างถูกต้อง สมเหตุสมผล ซึ่งนักวิจัยจะต้องตรวจสอบข้อมูลให้ทราบถึงลักษณะธรรมชาติของข้อมูลว่าตัวแปรมีลักษณะการแจกแจงแบบใด ความสัมพันธ์ระหว่างตัวแปรเป็นชนิดใด มีระดับการวัดอย่างไร ตัวแปรมีข้อมูลที่ขาดหาย (Missing) หรือไม่ ถ้าข้อมูลขาดหายมีจำนวนมากน้อยเท่าไร ข้อมูลมีค่าสุดโต่ง (Extremes) บ้างหรือไม่ และข้อมูลมีลักษณะไม่สอดคล้องกับข้อตกลงเบื้องต้น (Assumptions) ของสถิติที่จะใช้อย่างไรบ้าง (Schumacker & Lomax, 2010, p. 29) วิธีการตรวจสอบข้อมูลก่อนการดำเนินการวิเคราะห์ข้อมูลสามารถสรุปได้ดังนี้ (Kline, 2005, p. 48)

1. การตรวจสอบลักษณะการแจกแจงของตัวแปร (Distribution)

นักวิจัยต้องรู้ว่าตัวแปรใช้ในงานวิจัย มีระดับการวัดแบบใด การออกแบบการสุ่มตัวอย่าง และออกแบบเครื่องมือวิจัยรวบรวมข้อมูลให้มีคุณสมบัติตามที่ต้องการ การตรวจสอบคุณสมบัติที่ต้องทำเป็นอย่างแรก คือ การตรวจสอบลักษณะการแจกแจงของตัวแปร วิธีที่ง่ายที่สุดคือ การทำแผนภูมิฮิสโตแกรม (Histogram) สำหรับตัวแปรที่มีระดับวัดแบบนามบัญญัติ (Nominal Scale) และเรียงอันดับ (Ordinal Scale) และทำแผนภูมิต้น – ใบ (Stem and Leaf Plot) สำหรับตัวแปรที่มีระดับการวัดแบบอันตรภาค (Interval Scale) และอัตราส่วน (Ratio Scale) แผนภูมิที่ได้ของตัวแปรแต่ละตัวจะให้นักวิจัยทราบว่าลักษณะการแจกแจงของตัวแปรเป็นแบบใด นอกจากนี้อาจใช้การวิเคราะห์การแจกแจงความถี่ (Frequency Distribution) และการวิเคราะห์ด้วยสถิติบรรยาย ได้แก่ ค่าเฉลี่ย (Mean) มัชฐาน (Median) ฐานนิยม (Mode) พิสัย (Range) ส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation) ความเบ้ (Skewness) ความโค้ง (Kurtosis) ค่าสถิติเหล่านี้จะช่วยให้สามารถสรุปได้ว่าตัวแปรมีการแจกแจงเป็นโค้งปกติหรือไม่ โดยทั่วไปสถิติวิเคราะห์เกือบทุกชนิดที่เป็นสถิติขั้นสูงจะมีข้อตกลงเบื้องต้นว่า การแจกแจงของตัวแปรต้องเป็นแบบโค้งปกติ ซึ่งการตรวจสอบลักษณะการแจกแจงของตัวแปรต้องตรวจสอบทั้งแบบตัวแปรเดียว (Univariate) และตัวแปรพหุ (Multivariate) วิธีการตรวจสอบลักษณะการแจกแจงของตัวแปรเดียว

นอกจากวิธีการที่กล่าวมาในตอนต้นแล้ว ยังมีวิธีการตรวจสอบด้วยวิธีอื่น เช่น การทดสอบ
 เทียบความกลมกลืนด้วย Chi – square (Chi – square Goodness of Fit Test) การทดสอบ
 Kolmogorov – Smirnov และการทดสอบนัยสำคัญของความเบ้และความโด่งด้วยสถิติ Z ตามสูตร

$$Z_{\text{skewness}} = \frac{\text{Skewness}}{\sqrt{\frac{6}{n}}} \quad \text{และ} \quad Z_{\text{kurtosis}} = \frac{\text{Kurtosis}}{\sqrt{\frac{24}{n}}}$$

เมื่อ n คือ ขนาดกลุ่มตัวอย่าง ถ้าผลการทดสอบได้ค่า Z สูงกว่าค่าวิกฤติ จะปฏิเสธ
 สมมติฐานหลักที่แสดงว่าตัวแปรมีการแจกแจงไม่เป็นโค้งปกติ เพราะค่าความเบ้และความโด่ง
 แตกต่างจากค่าในโค้งปกติ คือ หากมีค่าเกินกว่า 2.58 แสดงว่าตัวแปรนั้นมีการแจกแจงที่เบี่ยงเบน
 จากโค้งปกติที่ระดับความเชื่อมั่น 99% หรือถ้ามีค่าเกินกว่า 1.96 แสดงว่าตัวแปรนั้นมีการแจกแจง
 ที่เบี่ยงเบนจากโค้งปกติที่ระดับความเชื่อมั่น 95% แต่ค่าคะแนนมาตรฐานของค่าความเบ้และ
 ความโด่งมีข้อด้อยในการใช้ เมื่อ n มีค่ามาก ๆ จะทำให้ค่ามาตรฐานมีค่าสูงทำให้เพิ่มโอกาส
 ในการสรุปว่าตัวแปรนั้นมีการแจกแจงไม่เป็นโค้งปกติ ทั้ง ๆ ที่ตัวแปรอาจแจกแจงเป็นปกติแล้ว
 ดังนั้นเมื่อ n มีจำนวนมาก จึงควรพิจารณาการแจกแจงของตัวแปร โดยดูรูปการแจกแจงมากกว่า
 การใช้สูตร ส่วนในการตรวจสอบตัวแปรพหุนั้นยังไม่มีเกณฑ์ที่แน่ชัดในการกำหนด
 แต่ส่วนมากมักจะให้ข้อสมมติ (Assume) ว่ามีการแจกแจงแบบปกติพหุ (Multivariate Normality)
 ถ้า 1) การแจกแจงในระดับตัวแปรเดียว (Univariate) ทุกตัวแปรมีการแจกแจงแบบปกติ
 2) การแจกแจงร่วมกันของตัวแปรแต่ละคู่มีการแจกแจงแบบปกติทวินาม (Bivariate Normality)
 และ 3) แผนภาพการกระจายของทุก ๆ ตัวแปรทวินาม (Bivariate) มีลักษณะเป็นเส้นตรง (Kline,
 2005, pp. 48 – 49) อีกวิธีหนึ่งคือตรวจสอบร่วมกับการวิเคราะห์เศษเหลือ (Residual Analysis)
 เช่น วิเคราะห์ตรวจสอบว่าเศษเหลือของสมการถดถอยพหุมีการแจกแจงแบบปกติหรือไม่
 เมื่อตรวจสอบว่าตัวแปรมีการแจกแจงไม่เป็นโค้งปกติ (Non – normal) วิธีการที่ควรจัดการ
 กับข้อมูล คือ การเปลี่ยนรูป (Transform) ข้อมูลตัวแปร เพื่อให้ได้การแจกแจงใกล้เคียงโค้งปกติ
 ถ้าการแจกแจงของตัวแปรมีความเบ้ทางลบให้เปลี่ยนรูปตัวแปรโดยใช้ค่ารากที่สอง ($\sqrt{x}$)
 ถ้าการแจกแจงของตัวแปรมีความเบ้ทางบวกให้เปลี่ยนรูปตัวแปรโดยใช้ค่าลอการิทึม ($\log x$)
 ถ้าการแจกแจงของตัวแปรมีความโด่งต่ำกว่าโค้งปกติให้เปลี่ยนรูปตัวแปรโดยใช้ส่วนกลับ
 ($1/x$) (Schumacker & Lomax, 2010, p. 36; Kline, 2005, pp. 50 – 51)

2. การตรวจสอบความสัมพันธ์ระหว่างตัวแปร

สถิติวิเคราะห์หาค่าความสัมพันธ์ส่วนใหญ่เป็นสถิติวิเคราะห์ที่มีพื้นฐานมาจากการวิเคราะห์การถดถอย และวิเคราะห์สหสัมพันธ์ ซึ่งมีข้อตกลงเบื้องต้นว่าความสัมพันธ์ระหว่างตัวแปรแต่ละคู่ต้องเป็นแบบเส้นตรง การตรวจสอบลักษณะความสัมพันธ์ระหว่างตัวแปรว่าเป็นแบบเส้นตรงหรือไม่ ทำได้หลายวิธี ได้แก่ การสร้างแผนภาพการกระจาย (Scatter Plot) ซึ่งสามารถสร้างแผนภูมิได้ทีละ 2 ตัวแปร หรือสร้างแผนภูมิเป็นเมตริกซ์ของแผนภูมิระหว่างตัวแปรทีละคู่หลายคู่พร้อมกันได้ หรือใช้วิธีการทดสอบสมมติฐานทางสถิติว่าลักษณะความสัมพันธ์เป็นแบบเส้นตรง (Linearity) หรือไม่ก็ได้ ในกรณีที่ความสัมพันธ์ระหว่างตัวแปรเป็นเส้นโค้ง นักวิจัยอาจเปลี่ยนรูปตัวแปร หรือเลือกใช้สถิติวิเคราะห์ที่เหมาะสม เช่น การวิเคราะห์ถดถอยแบบพหุนามเป็นต้น

3. การตรวจสอบข้อมูลขาดหาย (Missing)

ข้อมูลขาดหายเป็นเรื่องที่หลีกเลี่ยงไม่ได้ในการวิจัย เมื่อมีข้อมูลขาดหายสิ่งแรกที่นักวิจัยควรทำคือ การตรวจสอบว่าข้อมูลที่ขาดหายนั้นเป็นการขาดหายโดยสุ่ม (Random) หรือขาดหายแบบมีระบบ (Systematic) ถ้าหน่วยตัวอย่างที่ขาดหายไม่มีลักษณะร่วมกัน หมายความว่าข้อมูลที่ขาดหายนั้นเกิดขึ้นโดยสุ่ม นักวิจัยสามารถดำเนินการวิเคราะห์ข้อมูลต่อไปได้ ถ้าหน่วยตัวอย่างที่ขาดหายมีลักษณะคล้ายกัน นักวิจัยจะต้องเก็บข้อมูลเพิ่มเติม สำหรับวิธีการจัดการกับข้อมูลที่ขาดหายแบบง่ายทำได้ 3 วิธี คือ 1) ตัดข้อมูลส่วนที่ขาดหายเป็นคู่ (Pairwise Deletion) 2) ตัดข้อมูลส่วนที่ขาดหายของหน่วยตัวอย่างหน่วยนั้นทิ้งทั้งหมด (Listwise Deletion) และ 3) ใช้สถิติวิเคราะห์ประมาณค่าข้อมูลที่ขาดหายใส่แทน (Replacement of Missing Data) การประมาณค่าอาจใช้ค่าเฉลี่ยหรือค่าประมาณจากการวิเคราะห์ถดถอยก็ได้ (Schumacker & Lomax, 2010, p. 20)

4. การตรวจสอบข้อมูลสุดโต่ง (Extremes or Outliers)

ตัวแปรที่มีค่าของตัวแปรค่าหนึ่งค่าใดสูงมาก หรือต่ำมากผิดปกติ แสดงว่ามีข้อมูลสุดโต่ง วิธีการตรวจสอบข้อมูลสุดโต่งอาจทำไปพร้อมกับการตรวจสอบลักษณะการแจกแจงของตัวแปรและการตรวจสอบลักษณะความสัมพันธ์ระหว่างตัวแปร เพราะข้อมูลสุดโต่งมีได้ทั้งในการแจกแจงของตัวแปรเดียว (Univariate) แบบทวินาม (Bivariate) และแบบพหุนาม (Multivariate) นั่นคือ การวิเคราะห์ด้วยแผนภูมิต้น – ใบ และแผนภาพการกระจาย จะช่วยให้นักวิจัยค้นพบว่ามีข้อมูลสุดโต่งหรือไม่ (Schumacker & Lomax, 2010, p. 27) เมื่อพบแล้ว นักวิจัยต้องพิจารณาต่อไปว่าควรตัดหรือควรคงข้อมูลสุดโต่งไว้ ถ้าผลการวิเคราะห์ไม่แตกต่างกัน ก็ควรคงข้อมูลสุดโต่งไว้

5. การวิเคราะห์เพื่อตรวจสอบข้อตกลงเบื้องต้นของสถิติวิเคราะห์

ในการวิเคราะห์ข้อมูลด้วยสถิติขั้นสูง การตรวจสอบว่าข้อมูลสอดคล้องกับข้อตกลงเบื้องต้นของสถิติเป็นสิ่งจำเป็น เพราะการวิเคราะห์ข้อมูลที่มีตัวแปรหลายตัวนั้น ถ้าตัวแปร มีคุณสมบัติไม่สอดคล้องกับเบื้องต้นจะทำให้ผลการวิเคราะห์มีความเบี่ยงเบนและเอนเอียง ได้มากกว่าเมื่อมีตัวแปรน้อยตัว เหตุผลอีกประการหนึ่งคือ กระบวนการวิเคราะห์ข้อมูลด้วยสถิติ ขั้นสูงมีความยุ่งยากซับซ้อน เมื่อข้อมูลไม่เป็นไปตามข้อตกลงเบื้องต้น ผลการวิเคราะห์ข้อมูล ด้วยสถิติขั้นสูงจะแสดงผลการวิเคราะห์โดยพรางลักษณะที่ไม่สอดคล้องกับข้อตกลงเบื้องต้น ซึ่งถ้าใช้สถิติวิเคราะห์เบื้องต้นจะเห็นลักษณะที่ไม่สอดคล้องกับข้อตกลงเบื้องต้นได้ชัดเจน หากนักวิจัยไม่ตรวจสอบข้อมูลก่อน อาจจะได้ผลการวิเคราะห์ข้อมูลด้วยสถิติขั้นสูงที่มีอันตราย จากการที่ข้อมูลไม่เป็นไปตามข้อตกลงเบื้องต้น โดยที่นักวิจัยไม่สามารถสังเกตได้

ขั้นตอนการวิเคราะห์ข้อมูลด้วยโมเดลสมการเชิงโครงสร้าง

ขั้นที่ 1 การศึกษาทฤษฎีและงานวิจัยที่เกี่ยวข้อง (Review Literatures) ความสำคัญ ของการศึกษาทฤษฎีและงานวิจัยที่เกี่ยวข้องกับเรื่องที่ต้องการศึกษานอกจากจะทำให้ นักวิจัย สามารถพัฒนากรอบแนวคิดของการวิจัยได้เหมาะสมแล้ว ยังช่วยให้ นักวิจัยทราบว่าควรเลือก ตัวแปรใดบ้างเข้ามาอยู่ใน โมเดลและทำให้ทราบว่าตัวแปรที่เลือกมานั้นควรสร้างเครื่องมือใด ตัวแปรเหล่านั้นอย่างไร

ขั้นที่ 2 การกำหนดลักษณะเฉพาะของโมเดล (Model Specification) หลังจากที่ได้ศึกษา ทฤษฎีอย่างดีเพียงพอแล้วจะสามารถนำตัวแปรต่าง ๆ ที่เกี่ยวข้องกับการวิจัยมาพัฒนาเป็นกรอบ แนวคิดของการวิจัย โดยกำหนดความสัมพันธ์ระหว่างตัวแปรแฝงหรือตัวแปรสังเกตได้และ พารามิเตอร์ใน โมเดลให้เป็น โมเดลการวิจัยของนักวิจัย (Kline, 2005, p. 64) ซึ่ง Cooley (1978 cited in Schumacker & Lomax, 2010, p. 55) กล่าวว่า เป็นขั้นตอนที่ยากที่สุดในการวิเคราะห์โมเดล SEM เพราะต้องอาศัยการศึกษาทฤษฎีที่เข้มแข็งในการกำหนดโครงสร้างโมเดล

ขั้นที่ 3 การระบุความเป็นไปได้ค่าเดียวของโมเดล (Model Identification) เป็นการศึกษา ลักษณะการกำหนดค่าพารามิเตอร์ที่ยังไม่ทราบค่าใน โมเดลการวิจัยว่าเป็นไปตามเงื่อนไข ของการวิเคราะห์หรือไม่ หรือเป็นการระบุว่ามีโมเดลนั้นสามารถนำมาประมาณค่าพารามิเตอร์ได้ เป็นค่าเดียวหรือไม่ โดยการเปรียบเทียบค่า $p(p+1)/2$ กับจำนวนพารามิเตอร์ที่ต้องการประมาณค่า เมื่อ p คือ จำนวนตัวแปรสังเกตได้ทั้งหมดใน โมเดล โดยมีเงื่อนไขการพิจารณาดังนี้ (Schumacker & Lomax, 2010, p. 56 – 58)

1. ถ้า $p(p+1)/2$ มีค่าน้อยกว่าจำนวนพารามิเตอร์ที่ต้องการประมาณค่า จะเรียกโมเดล นั้นว่า โมเดลระบุความเป็นไปได้ค่าเดียวไม่พอ (Under Identification) ค่า df เป็นลบ ซึ่งโมเดล ประเภทนี้จะไม่สามารถประมาณค่าพารามิเตอร์ได้

2. ถ้า $p(p+1)/2$ มีค่าเท่ากับจำนวนพารามิเตอร์ที่ต้องการประมาณค่า จะเรียกโมเดลนั้นว่า โมเดลระบุความเป็นไปได้ค่าเดียวพอดี (Just Identification) ซึ่งค่า df เป็นศูนย์ ซึ่งโมเดลประเภทนี้จะประมาณค่าพารามิเตอร์ได้ไม่ดี

3. ถ้า $p(p+1)/2$ มีค่ามากกว่าจำนวนพารามิเตอร์ที่ต้องการประมาณค่า จะเรียกโมเดลนั้นว่า โมเดลระบุความเป็นไปได้ค่าเดียวเกินพอดี (Over Identification) ซึ่งค่า df เป็นบวก และโมเดลประเภทนี้จะสามารถประมาณค่าพารามิเตอร์ต่าง ๆ ในโมเดลได้

การระบุความเป็นไปได้ค่าเดียวทำให้นักวิจัยทราบได้ล่วงหน้าว่า โมเดลนั้นสามารถประมาณค่าพารามิเตอร์ได้หรือไม่ ดังนั้นนักวิจัยจึงควรระบุความเป็นไปได้ค่าเดียวของโมเดลก่อนดำเนินการวิเคราะห์ข้อมูล

ขั้นที่ 4 การประมาณค่าพารามิเตอร์ของโมเดล (Parameter Estimation of the Modal) เมื่อตรวจสอบว่าโมเดลระบุความเป็นไปได้ค่าเดียวเกินพอดีแล้ว ต้องทำการประมาณค่าพารามิเตอร์ทุกค่าในโมเดล ซึ่งหลักการประมาณค่าพารามิเตอร์ คือ ตรวจสอบความกลมกลืนระหว่างโมเดลสมมติฐานกับข้อมูลเชิงประจักษ์ การเปรียบเทียบใช้เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมเป็นตัวเกณฑ์ในการเปรียบเทียบ โดยนำเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมที่คำนวณได้จากกลุ่มตัวอย่างที่เป็นข้อมูลเชิงประจักษ์ (แทนด้วยสัญลักษณ์ S) มาเทียบกับเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมที่ถูกสร้างขึ้นจากพารามิเตอร์ที่ประมาณค่าได้จากโมเดลสมมติฐานวิจัย (แทนด้วยสัญลักษณ์ $\hat{\Sigma}$) ถ้าเมตริกซ์ทั้งสองมีค่าใกล้เคียงกัน หมายความว่า โมเดลสมมติฐานวิจัยมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ เนื่องจากเมตริกซ์ $\hat{\Sigma}$ ถูกสร้างจากค่าประมาณของพารามิเตอร์ ดังนั้นการประมาณค่าพารามิเตอร์จึงใช้หลักการวิเคราะห์เปรียบเทียบความกลมกลืนระหว่างเมตริกซ์ดังกล่าวเป็นเงื่อนไขในการประมาณค่าพารามิเตอร์ จุดมุ่งหมายของการประมาณค่าพารามิเตอร์ คือ การหาพารามิเตอร์ที่จะทำให้เมตริกซ์ S และ $\hat{\Sigma}$ มีค่าใกล้เคียงกันมากที่สุด การกำหนดเงื่อนไขให้เมตริกซ์ S และ $\hat{\Sigma}$ มีค่าใกล้เคียงกันนั้น ใช้การสร้างฟังก์ชันความกลมกลืน (Fit Function) เป็นตัวเกณฑ์ในการตรวจสอบ (Schumacker & Lomax, 2010, p. 59 – 60)

ขั้นตอนที่ 5 การตรวจสอบความกลมกลืนของโมเดลการวิจัยกับข้อมูลเชิงประจักษ์ (Model Fit) โดยนำเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมจากการประมาณค่าตามโมเดล (Computed Covariance Matrix: $\hat{\Sigma}$) ลบจากเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมจากข้อมูลกลุ่มตัวอย่าง (Sample Covariance Matrix: S) เรียกเมตริกซ์นี้ว่า เมตริกซ์ส่วนที่เหลือ (Residual Covariance Matrix) เพื่อทดสอบว่าเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมจากการประมาณค่าตามโมเดลแตกต่างจากเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วม

จากข้อมูลกลุ่มตัวอย่างหรือไม่ โดยตั้งสมมติฐานศูนย์ คือ $H_0: S = \hat{S}$ และสมมติฐานทางเลือก คือ $H_a: S \neq \hat{S}$ ถ้าค่าสถิติทดสอบที่ไม่มีนัยสำคัญจะแสดงว่าโมเดลการวิจัยกับข้อมูลเชิงประจักษ์ สอดคล้องกลมกลืนกัน นอกจากนี้ยังมีดัชนีแสดงความสอดคล้องกลมกลืนของโมเดลอีกหลายค่า (Schumacker & Lomax, 2010, p. 63)

ขั้นที่ 6 การปรับโมเดล (Model Modification) ถ้าโมเดลการวิจัยกับข้อมูลเชิงประจักษ์ ยังไม่สอดคล้องกลมกลืนกัน ผู้วิจัยจะต้องปรับโมเดล แล้วดำเนินการวิเคราะห์ใหม่จนกว่าโมเดล การวิจัยกับข้อมูลเชิงประจักษ์จะสอดคล้องกลมกลืนกัน จากนั้นจึงนำพารามิเตอร์ต่าง ๆ ในโมเดลไปเขียนรายงานได้ (Schumacker & Lomax, 2010, p. 64)

ประเภทของพารามิเตอร์ในโมเดลสมการเชิงโครงสร้าง

พารามิเตอร์ในโมเดลสมการเชิงโครงสร้างจำแนกได้เป็น 3 ประเภท คือ พารามิเตอร์ อิสระ (Free Parameter) พารามิเตอร์คงที่ (Fixed Parameter) และพารามิเตอร์บังคับ (Constrain Parameter) (Schumacker & Lomax, 2010, pp. 379 – 382)

1. พารามิเตอร์อิสระ (Free Parameter) หมายถึง พารามิเตอร์ที่ไม่ทราบค่า ซึ่งต้องการ ให้มีการประมาณค่า และมีได้บังคับให้มีค่าอย่างใดอย่างหนึ่ง ได้แก่ สัมประสิทธิ์ถดถอยในโมเดล โครงสร้าง หรือน้ำหนักองค์ประกอบใน โมเดลการวัด ที่ผู้วิจัยต้องการทราบค่า

2. พารามิเตอร์คงที่ (Fixed Parameter) มีได้ 2 แบบ ซึ่งแบบแรก คือ พารามิเตอร์ ที่ไม่ต้องการให้มีการประมาณค่า หรือให้มีค่าเป็นศูนย์ เพราะเป็นค่าพารามิเตอร์ที่รอบทฤษฎี หรือเอกสารงานวิจัยไม่ได้ระบุว่ามีความจำเป็นของพารามิเตอร์นี้ หรือบอกได้ว่าความสัมพันธ์ระหว่าง ตัวแปรมีค่าเป็นศูนย์ ในกรณีเช่นนี้ ผู้วิจัยต้องกำหนดให้พารามิเตอร์เหล่านี้มีค่าเป็นศูนย์หรือไม่ ต้องให้โปรแกรมประมาณค่าพารามิเตอร์เหล่านี้ แบบที่สอง คือ พารามิเตอร์ที่ไม่ทราบค่า แต่ต้องการประมาณค่าแล้วให้มีค่าเท่ากับตัวเลขค่าใดค่าหนึ่งที่ไม่เท่ากับศูนย์ซึ่งเป็นค่าใด ๆ ก็ตาม ที่ผู้วิจัยต้องการ

3. พารามิเตอร์บังคับ (Constrained Parameter) หมายถึง พารามิเตอร์ที่ไม่ทราบค่า แต่ต้องการให้โปรแกรมประมาณค่าให้เท่ากับพารามิเตอร์ตัวอื่นตามที่ผู้วิจัยได้ระบุไว้ให้มีค่า เท่ากัน ซึ่งจะใช้มากในการวิเคราะห์โมเดลกลุ่มพหุ (Multiple Group) หรือการทดสอบโมเดล ตั้งแต่ 2 โมเดลที่เหมือนกันตั้งแต่ 2 กลุ่มขึ้นไป

โมเดลการวิเคราะห์ห้้องค์ประกอบเชิงยืนยัน (Confirmatory Factor Analysis Model)

ความหมายของโมเดลการวิเคราะห์ห้้องค์ประกอบเชิงยืนยัน

โมเดลการวิเคราะห์ห้้องค์ประกอบเชิงยืนยัน (CFA Model) เป็นส่วนหนึ่งของ โมเดลสมการเชิงโครงสร้าง (SEM Model) คือ เป็นส่วนของ โมเดลการวัด (Measurement Model) ซึ่งเป็นโมเดลที่ระบุความสัมพันธ์เชิงเส้นระหว่างตัวแปรแฝงกับตัวแปรสังเกตได้ (Schumacker & Lomax, 2010, p. 184; Brown, 2006, p. 51) ซึ่งการวิเคราะห์ห้้องค์ประกอบเชิงยืนยัน (CFA) อยู่บนฐานของการวิเคราะห์ห้้องค์ประกอบเชิงสำรวจ (EFA) (Kline, 2005, pp. 204 – 205) ซึ่งมีวัตถุประสงค์ 3 ประการ คือ 1) เพื่อตรวจสอบทฤษฎีที่ใช้เป็นพื้นฐานในการวิเคราะห์ห้้องค์ประกอบ 2) เพื่อสำรวจและระบุห้้องค์ประกอบ และ 3) เพื่อเป็นเครื่องมือในการสร้างตัวแปรใหม่ การวิเคราะห์ห้้องค์ประกอบเชิงยืนยันเป็นการศึกษาทฤษฎีและงานวิจัยที่เกี่ยวข้องก่อนว่าคุณลักษณะที่ผู้วิจัยต้องการศึกษามีห้้องค์ประกอบอะไรบ้าง ห้้องค์ประกอบนั้น ๆ วัดได้ด้วยตัวแปรสังเกตได้ อะไรบ้าง จากนั้นกำหนดเป็น โมเดลห้้องค์ประกอบ แล้วเก็บข้อมูลตัวแปรสังเกตได้ต่าง ๆ ที่กำหนด แล้ววิเคราะห์ห้้องค์ประกอบที่กำหนดสอดคล้องกับข้อมูลเชิงประจักษ์หรือไม่ ซึ่งมีกระบวนการตรวจสอบความตรงของ โมเดลที่ชัดเจน ผลการวิเคราะห์ห้้องค์ประกอบให้ค่าพารามิเตอร์ รวมทั้งผลการทดสอบนัยสำคัญของพารามิเตอร์ การทดสอบความไม่แปรเปลี่ยนของ โมเดล ดังนั้น เทคนิค CFA จึงเป็นเครื่องมือที่นักวัดผลนำมาใช้ในการศึกษาและพัฒนาคุณภาพของเครื่องมือ เช่น แบบวัดแบบสอบถาม ได้เป็นอย่างดี (Brown, 2006, p. 1)

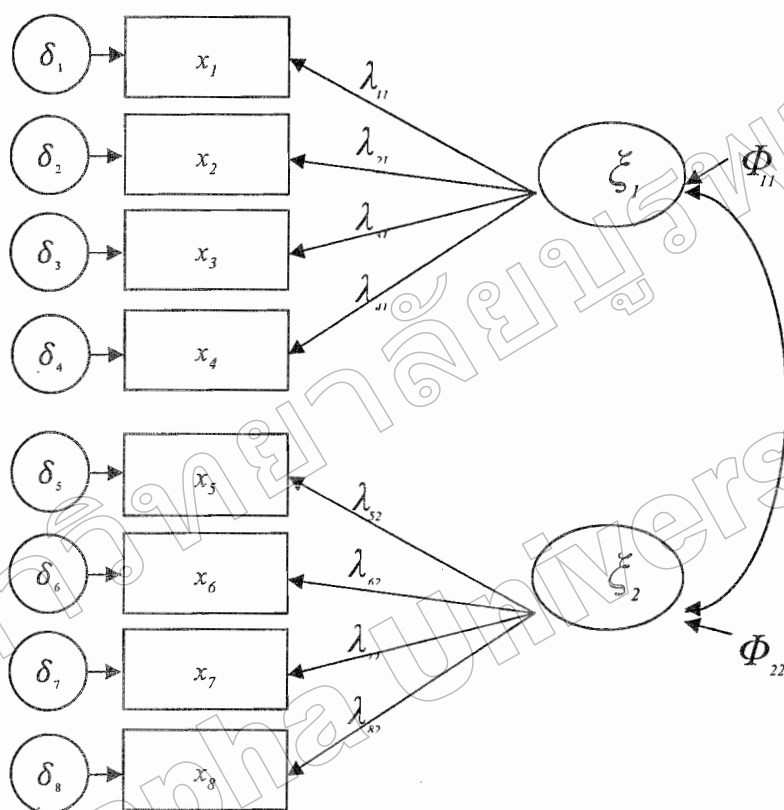
ขั้นตอนการวิเคราะห์ห้้องค์ประกอบเชิงยืนยัน (Schumacker & Lomax, 2010,

pp. 163–174)

1. ศึกษาทฤษฎี งานวิจัยที่เกี่ยวข้องเกี่ยวกับความสัมพันธ์ของตัวแปรแฝง และตัวแปรสังเกตได้ (Review Literatures)
2. กำหนดโมเดลเชิงทฤษฎี (Model Conceptualization)
3. กำหนดภาพโครงสร้างห้้องค์ประกอบ (Factor Diagram Construction)
4. กำหนดลักษณะเฉพาะของโมเดล (Model Specification)
5. ระบุความเป็นไปได้ค่าเดียวของโมเดล (Model Identification)
6. ประมาณค่าพารามิเตอร์ (Parameter Estimate)
7. ตรวจสอบความสอดคล้องของโมเดลกับข้อมูลเชิงประจักษ์ (Assessment of Model Fit)
8. ปรับโมเดล (Model Modification)

การกำหนดประเภทของพารามิเตอร์และเมตริกซ์ที่ใช้ในโมเดลการวิเคราะห์องค์ประกอบ
เชิงยืนยัน

การกำหนดประเภทของพารามิเตอร์ใน โมเดล CFA ตามโมเดลดังนี้



ภาพที่ 2-2 โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน (CFA Model) (Brown, 2006, p. 55)

เมื่อ x_i	แทน ตัวแปรสังเกตได้
λ_{ij}	แทน น้ำหนักองค์ประกอบของตัวแปรสังเกตได้ i บนตัวแปรแฝง j
ξ_j	แทน เวกเตอร์ตัวแปรแฝง
δ_i	แทน เวกเตอร์ความคลาดเคลื่อนในการวัดตัวแปรสังเกตได้
Φ_{ij}	แทน เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมระหว่างตัวแปรแฝง

จากภาพที่ 2-2 เป็นโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน 2 องค์ประกอบ
โดยเมตริกซ์ที่ใช้ในโมเดล ได้แก่

จากภาพที่ 2-2 เป็นโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันโดยตัวแปร $x_1 - x_4$ เป็นตัวแปรสังเกตได้ที่อธิบายได้ด้วยองค์ประกอบ ζ_1 โดยมีน้ำหนักองค์ประกอบเท่ากับ λ_{11} , λ_{12} , λ_{13} และ λ_{14} ตามลำดับ ตัวแปร $x_5 - x_8$ เป็นตัวแปรสังเกตได้ขององค์ประกอบ ζ_2 โดยมีน้ำหนักองค์ประกอบเท่ากับ λ_{21} , λ_{22} , λ_{23} และ λ_{24} ตามลำดับ โดยทั่วไปค่าน้ำหนักองค์ประกอบจะถูกอธิบายเป็นค่าสัมประสิทธิ์การถดถอย (Regression Coefficient) ในรูปค่ามาตรฐาน (Standardized) และรูปทั่วไป (Unstandardized) (Kline, 2005, p. 167) และตัวแปรแฝงทุก ๆ ตัวแปรจะต้องมีหน่วย (Kline, 2005, pp. 169 – 170) ดังนั้น การกำหนดให้ค่าในรูปทั่วไปของ $\lambda_{11} = 1.00$ และ $\lambda_{21} = 1.00$ เป็นพารามิเตอร์คงที่ (Fixed Parameter) ก็เพื่อกำหนดหน่วยการวัด หรือสเกล (Scale) ให้กับตัวแปรแฝงให้มีหน่วยเดียวกับตัวแปรสังเกตได้ตัวนั้นนั่นเอง ตามหลักของการวิเคราะห์องค์ประกอบ จะต้องประมาณค่าน้ำหนักองค์ประกอบ เพื่อนำไปเสนอผลการวิจัย ค่าน้ำหนักองค์ประกอบเหล่านี้จึงเป็นค่าพารามิเตอร์ที่ต้องการทราบค่า จึงต้องกำหนดเป็นพารามิเตอร์อิสระ นอกจากนี้ยังมีค่าพารามิเตอร์ที่ต้องการทราบค่าซึ่งเป็นพารามิเตอร์อิสระอีก คือ $\delta_1 - \delta_8$ (Unique Variance) คือ ความแปรปรวนของ (Measurement Error) ของตัวแปรสังเกตได้ $x_1 - x_8$ ตามลำดับ ในโมเดล CFA อนุญาตให้ความคลาดเคลื่อนของการวัดสามารถมีความสัมพันธ์กันได้ (Kline, 2005, p. 168) ϕ_{11} และ ϕ_{22} เป็นความแปรปรวนขององค์ประกอบ ζ_1 และ ζ_2 ตามลำดับ และ ϕ_{12} เป็นความแปรปรวนร่วมหรือความสัมพันธ์ขององค์ประกอบ ζ_1 และ ζ_2 หากพิจารณาจากภาพที่ 2-2 จะเห็นว่าไม่มีพารามิเตอร์ที่เป็นน้ำหนักองค์ประกอบจาก ζ_1 ไปยังตัวแปร $x_5 - x_8$ ดังนั้น พารามิเตอร์เช่นนี้ จึงเป็นพารามิเตอร์คงที่ (Fixed) คือ ถือว่าน้ำหนักองค์ประกอบเหล่านี้มีค่าเป็นศูนย์ เพราะทฤษฎีไม่ได้บอกว่าตัวแปร $x_5 - x_8$ เป็นตัวแปรที่ใช้วัดองค์ประกอบ ζ_1 ได้ (Brown, 2006, p. 49)

Kline (2005, p. 172) กล่าวถึงเงื่อนไขของจำนวนตัวแปรสังเกตได้/ ตัวบ่งชี้ (Indicators) ในโมเดล CFA ว่า สำหรับโมเดล CFA ที่มีองค์ประกอบเดียวควรมีอย่างน้อย 3 ตัวบ่งชี้ และโมเดลที่มีตั้งแต่ 2 องค์ประกอบขึ้นไป สามารถมีอย่างน้อย 2 ตัวบ่งชี้ต่อองค์ประกอบ ซึ่ง Bollen (1989 cited in Kline, 2005, p. 172) ให้ข้อเสนอแนะว่าในโมเดลที่มีตั้งแต่ 2 องค์ประกอบขึ้นไป จำนวน 2 ตัวบ่งชี้ต่อองค์ประกอบมีแนวโน้มที่จะเกิดปัญหาในการประมาณค่า ดังนั้น ควรมีอย่างน้อย 3 ตัวบ่งชี้ต่อองค์ประกอบ โดยเฉพาะในกลุ่มตัวอย่างขนาดเล็ก ถ้าโมเดลมีตัวบ่งชี้มากขึ้นควรเพิ่มขนาดกลุ่มตัวอย่างมากขึ้นด้วยเช่นกัน

ผลลัพธ์จากการประมาณค่าในโมเดล CFA ความแปรปรวนร่วมระหว่างองค์ประกอบ ค่าน้ำหนักองค์ประกอบของตัวแปรสังเกตได้บนตัวแปรแฝง และค่าความคลาดเคลื่อนจากการวัดของแต่ละตัวแปรสังเกตได้ หากนักวิจัยมีการกำหนดลักษณะเฉพาะของโมเดลอย่างถูกต้องแล้ว

จะทำให้ผลลัพธ์การประมาณค่าต่าง ๆ ได้ผลดังนี้ คือ 1) ตัวแปรสังเกตได้ที่ใช้วัดองค์ประกอบเดียวกันมีความสัมพันธ์กันสูงและค่าน้ำหนักองค์ประกอบของตัวแปรสังเกตได้มีค่าสูง (มากกว่า .60) (Marsh & Hau, 1999 cited in Kline, 2005, p. 178) และ 2) ค่าความสัมพันธ์ระหว่างองค์ประกอบมีค่าไม่สูงเกินไป (ไม่เกิน .85) (Kline, 2005, p. 73)

ขนาดกลุ่มตัวอย่างที่ใช้ในการวิเคราะห์องค์ประกอบเชิงยืนยัน และโมเดลสมการเชิงโครงสร้าง

โดยทั่วไปในการวิเคราะห์สถิติประเภทพหุตัวแปร (Multivariate Statistics) มีหลักการกำหนดขนาดกลุ่มตัวอย่าง คือ การเป็นตัวแทนที่ดีของประชากร ในขณะเดียวกันก็ต้องคำนึงถึงเทคนิคในการประมาณค่าและข้อตกลงเบื้องต้น แต่อย่างไรก็ตามในทางปฏิบัตินักวิจัยยังต้องการใช้กลุ่มตัวอย่างที่น้อยที่สุดเท่าที่จะเป็นไปได้ แม้ว่าจะทราบว่าการใช้กลุ่มตัวอย่างขนาดใหญ่จะมีโอกาสที่ตัวแปรมีการแจกแจงเป็นโค้งปกติมากกว่ากลุ่มตัวอย่างที่เล็กกว่า และให้ค่าสถิติที่ค่อนข้างคงที่กว่า ข้อเสนอแนะเกี่ยวกับขนาดกลุ่มตัวอย่างที่ใช้ในงานวิจัยเกี่ยวกับโมเดลสมการเชิงโครงสร้างจากการศึกษาต่าง ๆ ดังนี้

1. กำหนดขนาดกลุ่มตัวอย่างจากจำนวนตัวแปรแฝงใน โมเดลและค่า Communalities ของตัวแปรสังเกตได้

Hair et al. (2010, p. 662) ได้เสนอขนาดกลุ่มตัวอย่างที่น้อยที่สุดในการวิเคราะห์โมเดล SEM ภายใต้งื่อนไขต่าง ๆ คือ 1) จำนวนตัวแปรแฝงในโมเดล 2) ค่า Communalities ของตัวแปรสังเกตได้ ซึ่งหมายถึง ปริมาณความแปรปรวนที่ชุดตัวแปรสังเกตได้อธิบายโมเดลการวัด และ 3) ปัญหาการระบุลักษณะเฉพาะของโมเดล ซึ่งแสดงได้ดังตารางที่ 2 – 2

ตารางที่ 2 – 2 ขนาดกลุ่มตัวอย่างที่น้อยที่สุดในการวิเคราะห์โมเดล SEM

จำนวนตัวแปรแฝง	ค่า Communalities	ขนาดกลุ่มตัวอย่างที่น้อยที่สุด
1. ไม่เกิน 5 ตัวแปร แต่ละตัวแปร มีมากกว่า 3 ตัวแปรสังเกตได้	> .60	100
2. ไม่เกิน 7 ตัวแปร	ประมาณ .50	150
3. ไม่เกิน 7 ตัวแปร	< .45	300
4. มีตัวแปรจำนวนมาก	ต่ำ	500

Hair et al. (2010, p. 662) กล่าวว่า การเพิ่มขนาดกลุ่มตัวอย่างจะทำเมื่อ 1) ข้อมูลมีการเบี่ยงเบนจากการแจกแจงแบบปกติพหุ (Multivariate Normal Distribution) 2) มีการใช้วิธีการประมาณค่าบางชนิดที่ต้องใช้กลุ่มตัวอย่างขนาดใหญ่ เช่น วิธี WLS/ ADF และ 3) มีข้อมูลสูญหายมากกว่า 10% แต่ถ้าสูญหายเกิน 15% ไม่ควรใช้เทคนิค SEM

2. กำหนดขนาดกลุ่มตัวอย่างจากจำนวนตัวแปรสังเกตได้ในโมเดล (Subject – to – variables Ratio)

Bentler and Chou (1987 cited in Schumacker & Lomax, 2010, p. 42) แนะนำว่าสัดส่วนระหว่างจำนวนตัวอย่าง (n) ต่อตัวแปรสังเกตได้ในโมเดล (p) คือ $n:p = 5:1$ เป็นขนาดเล็กที่สุดที่เพียงพอสำหรับตัวแปรที่มีการแจกแจงเป็นโค้งปกติ หรือ $n:p = 10:1$ สำหรับตัวแปรที่มีการแจกแจงแบบอื่น ๆ

Raykov and Marcoulides (2006 cited in Chumney, 2012, p. 29) กล่าวถึง Rule of Thumb ว่าขนาดกลุ่มตัวอย่างที่น้อยที่สุดสำหรับการประมาณค่าในโมเดลสมการเชิงโครงสร้างคือ 10 เท่าของจำนวนพารามิเตอร์อิสระในโมเดล

Costello and Osborne (2005 cited in Schumacker & Lomax, 2010, p. 42) ได้ศึกษาการจำลองข้อมูลด้วยวิธีการมอนติคาร์โล พบว่า ขนาดกลุ่มตัวอย่างที่มีความเหมาะสมที่สุดสำหรับการวิเคราะห์องค์ประกอบ คือ ขนาด 20 เท่า ของตัวแปรสังเกตได้ในโมเดล สอดคล้องกับข้อเสนอแนะของ Kline (2005, p. 178)

3. กำหนดขนาดกลุ่มตัวอย่างโดยใช้สูตรของ Cochran (1977 cited in Bartlett, Kotrlik, & Higgins, 2001, pp. 43 – 50)

Bartlett, Kotrlik and Higgins (2001, pp. 43 – 49) ได้ศึกษาการกำหนดขนาดกลุ่มตัวอย่างในการวิเคราะห์ถดถอยเชิงพหุ (Multiple Regression) และการวิเคราะห์องค์ประกอบ (Factor Analysis) สำหรับตัวแปรแบบต่อเนื่อง (Continuous Variables) และตัวแปรแบบจำแนกกลุ่ม (Categorical Variables) โดยใช้สูตรของ Cochran (1977 cited in Bartlett, Kotrlik, & Higgins, 2001, p. 48) โดยเปรียบเทียบในค่าแอลฟา (Alpha Level) ที่แตกต่างกัน 3 ระดับ ซึ่งสูตรของ Cochran นั้นผู้วิจัยต้องคำนึงถึง 2 ประการ คือ 1) ค่าความคลาดเคลื่อนที่สามารถยอมรับได้ในการวิจัย ซึ่งเรียกว่า Margin of Error และ 2) ระดับค่าแอลฟา (Alpha Level) คือระดับที่ผู้วิจัยสามารถยอมรับได้ว่าค่าขอบเขตที่แท้จริงของความคลาดเคลื่อน (True Margin of Error) หรือ เรียกว่า ค่าความคลาดเคลื่อนประเภทที่ 1 (Type 1 Error)

ระดับค่าแอลฟา (Alpha Level) ที่ใช้ในการศึกษาส่วนใหญ่ถูกกำหนดที่ระดับ .05 และ .01 (Ary, Jacobs, & Razavieh, 1996 cited in Bartlett, Kotrlik, & Higgins, 2001, p. 45)

ส่วนค่าความคลาดเคลื่อนที่สามารถยอมรับได้ (Margin of Error) ส่วนใหญ่ในการวิจัยทางการศึกษา และสังคมศาสตร์นิยมกำหนดที่ 3% สำหรับข้อมูลแบบต่อเนื่อง และ 5% สำหรับข้อมูลแบบจำแนกกลุ่ม (Krejcie & Morgan, 1970 cited in Bartlett, Kotrlik, & Higgins, 2001, p. 45)

3.1 กำหนดขนาดกลุ่มตัวอย่างสำหรับข้อมูลแบบต่อเนื่อง (Continuous Data)

ในงานวิจัยที่ข้อมูลได้จากการวัดโดยมาตรวัดแบบจำแนกกลุ่มหลาย ๆ กลุ่ม ผู้วิจัยควรใช้มาตรวัดในหลายระดับจึงจะสามารถกำหนด (Assume) ว่ามีลักษณะต่อเนื่องได้ โดยคำนวณได้ดังสูตรต่อไปนี้ (Bartlett, Kotrlik, & Higgins, 2001, p. 46)

$$n = \frac{n_0}{1 + (n_0 / N)} \quad \text{เมื่อ} \quad n_0 = \frac{t^2 \times s^2}{d^2}$$

เมื่อ n คือ ขนาดกลุ่มตัวอย่าง

N คือ จำนวนประชากร

t คือ ค่าสถิติที่ที่กำหนดโดยค่าแอลฟา (เช่น $\alpha = .05, t = 1.96$)

s คือ อัตราส่วนระหว่าง จำนวนของมาตรวัด (Number of Point on Scale) กับจำนวนส่วนเบี่ยงเบนมาตรฐาน (ประมาณ 98% ของจำนวนมาตรวัด)

d คือ ผลคูณระหว่างจำนวนของมาตรวัดกับค่าความคลาดเคลื่อนที่สามารถยอมรับได้ (Margin of Error)

3.2 กำหนดขนาดกลุ่มตัวอย่างสำหรับข้อมูลแบบจำแนกกลุ่ม (Categorical Data)

คำนวณได้ดังสูตรต่อไปนี้ (Bartlett, Kotrlik, & Higgins, 2001, p. 47)

$$n = \frac{n_0}{1 + (n_0 / N)} \quad \text{เมื่อ} \quad n_0 = \frac{t^2 \times p^2}{c^2}$$

เมื่อ n คือ ขนาดกลุ่มตัวอย่าง

N คือ จำนวนประชากร

t คือ ค่าสถิติที่ที่กำหนดโดยค่าแอลฟา (เช่น $\alpha = .05, t = 1.96$)

p คือ ค่า Estimated the Standard Deviation of the Scale ซึ่งกำหนดเท่ากับ .50

c คือ ค่าความคลาดเคลื่อนที่สามารถยอมรับได้ (Margin of Error)

ผลการศึกษาแสดงสรุปการคำนวณจากสูตร ของ Cochran ได้ดังตารางที่ 2 – 3

ตารางที่ 2 – 3 ขนาดกลุ่มตัวอย่างในการวิเคราะห์องค์ประกอบสำหรับตัวแปรแบบต่อเนื่องและ
ตัวแปรจำแนกกลุ่ม ตามสูตรของ Cochran (Bartlett, Kotrlik, & Higgins, 2001, p. 48)

Population Size (N)	Sample Size (n)					
	Continuous Data (Margin of Error = .03)			Categorical Data (Margin of Error = .05)		
	$\alpha = .10$	$\alpha = .05$	$\alpha = .01$	$\alpha = .10$	$\alpha = .05$	$\alpha = .01$
100	46	55	68	74	80	87
200	59	75	102	116	132	154
300	65	85	123	143	169	207
400	69	92	137	162	196	250
500	72	96	147	176	218	286
600	73	100	155	187	235	316
700	75	102	161	196	249	341
800	76	104	166	203	260	363
900	76	105	170	209	270	382
1,000	77	106	173	213	278	399
1,500	79	110	183	230	306	461
2,000	83	112	189	239	323	499
4,000	83	119	198	254	351	570
6,000	83	119	209	259	362	598
8,000	83	119	209	262	367	613
10,000	83	119	209	264	370	623

อย่างไรก็ตาม Bartlett, Kotrlik and Higgins (2001, p. 49) ได้เสนอแนะสำหรับขนาด
กลุ่มตัวอย่างในการวิเคราะห์องค์ประกอบว่า กลุ่มตัวอย่างไม่ควรมีขนาดต่ำกว่า 100
เพราะเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นจะส่งผลให้ระดับนัยสำคัญ (ค่าแอลฟา) ของค่าน้ำหนัก
องค์ประกอบมีค่าน้อยลงได้

Marsh, Hau, Balla and Grayson (1998 cited in Myers, Ahn, & Jin, 2011, p. 412)
กล่าวถึง Rules of Thumb เกี่ยวกับขนาดกลุ่มตัวอย่างซึ่งเพียงพอที่ทำให้การประมาณค่า
สำหรับโมเดล CFA ให้ผลลัพธ์ที่มีความเหมาะสม คือ โมเดลมีความลู่เข้าค่าเดียว (Convergence)

ความแม่นยำของค่าสถิติ (Statistical Precision) และความมีประสิทธิภาพของค่าสถิติ (Statistical Power) ควรมีขนาดกลุ่มตัวอย่าง (n) ≥ 200

แต่อย่างไรก็ตามก็ยังไม่มีการตายตัวสำหรับขนาดตัวอย่างที่เพียงพอสำหรับ โมเดล ภายใต้งื่อนไขต่าง ๆ (Gagné & Hancock, 2006; Jackson, 2001; 2003 cited in Myers, Ahn, & Jin, 2011, p. 412) ปัญหาที่สำคัญของ Rules of Thumb เกี่ยวกับขนาดกลุ่มตัวอย่างที่เพียงพอใน โมเดล CFA นั้นขึ้นอยู่กับเงื่อนไขต่าง ๆ เช่น โมเดลไม่มีทฤษฎีสันับสนุนที่เข้มแข็ง การแจกแจงของตัวแปร ค่าเชื่อมั่นของตัวแปรสังเกตได้ ขนาดของโมเดล หรือการกำหนดลักษณะเฉพาะของโมเดลที่ไม่เหมาะสม

วิธีการประมาณค่าพารามิเตอร์ (Parameter Estimation Method) ที่ใช้ในการวิจัย

หลักการวิเคราะห์หาค่าประมาณเชิงยืนยัน คือ ตรวจสอบความกลมกลืนระหว่างโมเดล สมมติฐานกับข้อมูลเชิงประจักษ์ การเปรียบเทียบใช้เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมเป็นตัวเกณฑ์ในการเปรียบเทียบ โดยนำเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมที่คำนวณได้จากกลุ่มตัวอย่างที่เป็นข้อมูลเชิงประจักษ์ (Sample Covariance Matrix แทนด้วยสัญลักษณ์ S) มาเทียบกับเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมที่ถูกสร้างขึ้นจากพารามิเตอร์ที่ประมาณค่าได้จากโมเดลสมมติฐานวิจัย (Population Covariance Matrix แทนด้วยสัญลักษณ์ Σ) ถ้าเมตริกซ์ทั้งสองมีค่าใกล้เคียงกัน หมายความว่า โมเดลสมมติฐานวิจัยมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ เนื่องจากเมตริกซ์ Σ ถูกสร้างจากค่าประมาณของพารามิเตอร์ ดังนั้น การประมาณค่าพารามิเตอร์จึงใช้หลักการวิเคราะห์เปรียบเทียบความกลมกลืนระหว่างเมตริกซ์ดังกล่าวเป็นเงื่อนไขในการประมาณค่าพารามิเตอร์ จุดมุ่งหมายของการประมาณค่าพารามิเตอร์ คือ การหาค่าพารามิเตอร์ที่จะทำให้เมตริกซ์ S และ Σ มีค่าใกล้เคียงกันมากที่สุด (Flora & Curran, 2004, pp. 466 – 467)

การกำหนดเงื่อนไขให้เมตริกซ์ S และ Σ มีค่าใกล้เคียงกันนั้น ใช้การสร้างฟังก์ชันความกลมกลืน (Fit Function) เป็นตัวเกณฑ์ในการตรวจสอบ รูปแบบของฟังก์ชันความกลมกลืนที่ง่ายที่สุด คือ รูปแบบของผลต่างระหว่างเมตริกซ์ทั้งสอง แต่ฟังก์ชันลักษณะนี้ไม่สะดวกต่อการคำนวณเพื่อประมาณค่าพารามิเตอร์ นักสถิติจึงได้กำหนดรูปแบบของฟังก์ชันความกลมกลืนขึ้นหลายรูปแบบ อันเป็นที่มาของวิธีการประมาณค่าที่แตกต่างกันไป รูปแบบของฟังก์ชันที่กำหนดขึ้นนี้ ทุกฟังก์ชันต้องมีคุณสมบัติรวม 4 ประการต่อไปนี้ ซึ่งจะทำให้ได้ค่าประมาณที่มีความคงเส้นคงวา (Consistency) (Bollen, 1989, p. 106)

1. ฟังก์ชันความกลมกลืนต้องเป็นสเกลาร์ (Scalar) หรือเป็นเลขจำนวน
2. ฟังก์ชันความกลมกลืนต้องมีค่ามากกว่าหรือเท่ากับศูนย์
3. ฟังก์ชันความกลมกลืนต้องมีค่าเป็นศูนย์เมื่อเมตริกซ์ S และ $\hat{\Sigma}$ มีค่าเท่ากันเท่านั้น
4. ฟังก์ชันความกลมกลืนเป็นฟังก์ชันต่อเนื่อง (Continuous Function)

วิธีประมาณค่าพารามิเตอร์แบบต่าง ๆ เป็นการประมาณค่าที่ใช้ฟังก์ชันความกลมกลืนที่ต่างกัน ซึ่งผลจากการประมาณค่ามีคุณสมบัติของค่าประมาณแตกต่างกัน วิธีการประมาณค่าพารามิเตอร์นั้นจะต้องเป็นไปตามข้อตกลงเบื้องต้นจึงจะทำให้ผลการวิจัยนั้นน่าเชื่อถือซึ่งผลลัพธ์ที่ต้องการศึกษา ได้แก่ ค่าพารามิเตอร์ (Parameter Values) ค่าความคลาดเคลื่อนมาตรฐาน (Standard Error) และ ดัชนีวัดความสอดคล้องกลมกลืน (Fit Indices) รายละเอียดของการประมาณค่าพารามิเตอร์แต่ละวิธีสรุปได้ดังนี้ (Bollen, 1989, pp. 104 – 121)

1. วิธี ML (Maximum Likelihood)

การประมาณค่าพารามิเตอร์ด้วยวิธี ML เป็นเทคนิคที่ใช้ในการวิเคราะห์โมเดล SEM ที่มีความนิยมและใช้อย่างแพร่หลายมากที่สุด (Brown, 2006, pp. 72 – 73; Curran, West, & Finch, 1996, p. 17; Flora & Curran, 2004, p. 466; Kline, 2005, p. 178; Finney & DiStefano, 2006, p. 271) เป็นวิธีการที่อาศัยความน่าจะเป็นสูงสุดที่จะประมาณค่าพารามิเตอร์ให้ใกล้เคียงมากที่สุดจากข้อมูลเชิงประจักษ์จากกลุ่มตัวอย่าง การประมาณค่าพารามิเตอร์ในโมเดล CFA ใช้กระบวนการคำนวณทวนซ้ำ เพื่อจะได้ฟังก์ชันความกลมกลืนมีค่าน้อยที่สุด (Brown, 2006, p. 73) และเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วม จะใช้ฐานการคำนวณเมตริกซ์สหสัมพันธ์แบบ Pearson Product Moment (PPM) จากข้อมูลดิบมีฟังก์ชันความกลมกลืน ดังสูตร

$$F = \log |\hat{\Sigma}| + tr(S\hat{\Sigma}^{-1}) - \log |S| + k$$

เมื่อ tr คือ ผลรวมสมาชิกในแนวทแยงของเมตริกซ์

S คือ เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมจากกลุ่มตัวอย่าง

$\hat{\Sigma}$ คือ เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมที่ได้จากค่าประมาณพารามิเตอร์

k คือ จำนวนตัวแปรสังเกตได้ทั้งหมดในโมเดล

การประมาณค่าพารามิเตอร์ด้วยวิธี ML ซึ่งมีข้อตกลงเบื้องต้น ดังต่อไปนี้ (Flora & Curran, 2004, p. 467; Finney & DiStefano, 2006, p. 271; Brown, 2006, p. 73)

1. หน่วยตัวอย่างที่สังเกตต้องเป็นอิสระต่อกัน (Independent Observations)
2. กลุ่มตัวอย่างต้องมีขนาดใหญ่เพียงพอ (Large Sample Size)

3. มีการกำหนดลักษณะเฉพาะของโมเดลอย่างถูกต้อง หมายถึง การกำหนดโมเดลที่ใช้ประมาณค่าจะต้องตรงกับโครงสร้างจริงของประชากร (Proper Model Specification)
4. ข้อมูลจากตัวแปรสังเกตได้มีการแจกแจงปกติพหุ (Multivariate Normal Distribution)
5. ตัวแปรที่มีลักษณะเป็นข้อมูลแบบต่อเนื่อง (Continuous Data)

วิธีนี้ใช้ฟังก์ชันความถ่วงถ่วงที่ไม่ใช่ฟังก์ชันแบบเส้นตรง แต่ก็เป็นฟังก์ชันที่บอกความแตกต่างระหว่างเมตริกซ์ S และเมตริกซ์ $\hat{\Sigma}$ ได้ เพราะถ้าเมตริกซ์ทั้งสองมีค่าใกล้เคียงกัน พจน์แรกของฟังก์ชันจะมีค่าเท่ากับพจน์ที่สาม ในขณะที่พจน์กลางมีค่าเท่ากับศูนย์ ถ้าข้อมูลเป็นไปตามข้อตกลงเบื้องต้นดังกล่าวแล้ว ค่าประมาณของพารามิเตอร์ที่ได้จากวิธี ML จะมีคุณสมบัติดังนี้ คือ มีความคงเส้นคงวา มีประสิทธิภาพและเป็นอิสระจากมาตรวัดการแจกแจงสุ่มของค่าประมาณพารามิเตอร์ที่ได้จากวิธี ML เป็นแบบโค้งปกติ และความแปรปรวนของค่าประมาณขึ้นอยู่กับขนาดของค่าพารามิเตอร์

วิธี ML จะประมาณค่าได้ดีเมื่อข้อมูลมีการแจกแจงปกติพหุ (Multivariate Normal Distribution) และมีกลุ่มตัวอย่างที่มากพอ หากข้อมูลไม่เป็นไปตามข้อตกลงนี้ ค่าประมาณของพารามิเตอร์ ค่าความคลาดเคลื่อนมาตรฐานของค่าประมาณที่ได้จะมีความเอนเอียง และค่าสถิติ Chi-square ที่ใช้ทดสอบความสอดคล้องกลไกของโมเดลจะมีค่าสูงขึ้น ซึ่งส่งผลให้ผลลัพธ์จากการประมาณค่าพารามิเตอร์มีค่าไม่ถูกต้อง (Finney & DiStefano, 2006, p. 272; Brown, 2006, p. 73)

2. วิธี Robust WLS (Robust Weighted Least Square)

เป็นวิธีการประมาณค่าพารามิเตอร์ที่ปรับปรุงมาจากวิธี WLS (Weighted Least Square) หรือวิธี ADF (Asymptotically Distribution Free) โดยวิธีการ WLS/ ADF มีการประยุกต์ลงในซอฟต์แวร์ เช่น ในโปรแกรม LISREL และ Mplus ซึ่งทั้งสองโปรแกรมนี้จะให้ค่าพารามิเตอร์ ค่าเบี่ยงเบนมาตรฐาน และดัชนีวัดระดับความถ่วงถ่วงที่ใกล้เคียงกัน (Finney & DiStefano, 2006, p. 281)

Brown (1984 cited in Finney & DiStefano, 2006, p. 278) ได้พัฒนาการประมาณค่าด้วยเทคนิค WLS/ ADF โดยการปรับปรุงข้อบกพร่องของวิธี ML เกี่ยวกับข้อตกลงเบื้องต้นเรื่องการแจกแจงแบบปกติพหุและความแปรปรวนต่อการแจกแจงข้อมูลที่ไม่เป็นโค้งปกติ ค่าสถิติ Chi-square และค่าเบี่ยงเบนมาตรฐานที่ได้จากการประมาณค่าจึงมีความแปรปรวนต่อการแจกแจงข้อมูลที่ไม่เป็นโค้งปกติ มีฟังก์ชันความถ่วงถ่วง ดังสูตร

$$F = (s - \sigma)' W^{-1} (s - \sigma) = \sum \sum \sum \sum (w^{gh,ij} (s_{gh} - \sigma_{gh}) (s_{ij} - \sigma_{ij}))$$

เมื่อ s คือ สมาชิกในแนวทแยง และได้แนวทแยงของเมตริกซ์ S

σ คือ สมาชิกในแนวทแยง และได้แนวทแยงของเมตริกซ์ $\hat{\Sigma}$

W คือ เมตริกซ์ที่ใช้ถ่วงน้ำหนัก

$w^{gh,ij}$ คือ สมาชิกในแนวทแยง และได้แนวทแยงของอินเวอร์สของเมตริกซ์ W

การประมาณค่าพารามิเตอร์ด้วยวิธี WLS/ ADF เป็นวิธีที่มีการวางนัยทั่วไป

อย่างกว้างขวาง กล่าวได้ว่าวิธี ULS, GLS และ ML เป็นกรณีหนึ่งของวิธี WLS การประมาณค่าวิธีนี้ไม่ได้ใช้เมตริกซ์เต็มรูป แต่ใช้เฉพาะสมาชิกในแนวทแยง และใช้เมตริกซ์ W เป็นเมตริกซ์ถ่วงน้ำหนัก โดยถ่วงน้ำหนักด้วยอินเวอร์สของเมตริกซ์ W จุดด้อยของวิธี WLS คือ เมตริกซ์ W มีขนาดใหญ่มาก เช่น ถ้ามีตัวแปรสังเกตได้จำนวน 5 ตัวแปร เมตริกซ์ S มีขนาด 5×5 มีจำนวนสมาชิกเท่ากับ $(1/2)(5)(5+1)$ หรือ 15 จำนวน แต่เมตริกซ์ W มีขนาด 15×15 ตามจำนวนสมาชิกในเมตริกซ์ S ซึ่งทำให้โปรแกรมต้องใช้เวลามากในการประมาณค่า ข้อด้อยอีกประการหนึ่งคือ วิธีนี้ไม่เหมาะกับเมตริกซ์ที่มีการตัดข้อมูลสูญหาย (Missing) แบบตัดเฉพาะคู่ที่ขาด (Pairwise) ส่วนคุณสมบัติของพารามิเตอร์เหมือนกับวิธี ML

วิธีประมาณค่าแบบ WLS/ ADF มีความจำเป็นต้องใช้กลุ่มตัวอย่างขนาดใหญ่ เพราะทำให้การประมาณค่ามีความแม่นยำและมีความคลาดเคลื่อนเข้าหาค่าพารามิเตอร์จริงมากขึ้น ซึ่ง Jöreskog and Sörbom (1996 cited in Finney & DiStefano, 2006, p. 279) ได้แนะนำว่าขนาดกลุ่มตัวอย่างที่เล็กที่สุดควรมีค่าไม่ต่ำกว่า $1.5 (p)(p + 1)$ เมื่อ p คือ จำนวนตัวแปรสังเกตได้ในโมเดล CFA ในทางทฤษฎีแล้ว เมื่อโมเดลมีการแจกแจงข้อมูลไม่เป็นโค้งปกติ การประมาณค่าพารามิเตอร์ด้วยวิธีการ WLS/ ADF ควรจะให้การประมาณค่าพารามิเตอร์และให้ค่าดัชนีวัดความสอดคล้องกลมกลืนที่มีคุณสมบัติตามที่คาดหวัง (Browne, 1984 cited in Finney & DiStefano, 2006, p. 280) แต่อย่างไรก็ตามวิธีการ WLS/ ADF ก็มีจุดด้อย เช่น ภายใต้งานไขโมเดลมีขนาดปานกลางถึงขนาดใหญ่ (มีองค์ประกอบมากกว่า 2 องค์ประกอบ หรือมีตัวแปรสังเกตได้มากกว่า 8 ตัวแปร) หรือ กลุ่มตัวอย่างขนาดเล็กถึงปานกลาง ($n < 500$) โดยเฉพาะอย่างยิ่งในกรณีที่ข้อมูลมีลักษณะเป็นแบบไม่ต่อเนื่องที่มีการแจกแจงไม่เป็นโค้งปกติ

จากข้อจำกัดดังกล่าว Muthén (1984 cited in Finney & DiStefano, 2006, p. 292)

จึงได้พัฒนาวิธีการประมาณค่าแบบ WLS/ ADF ที่มีความแกร่งมากขึ้น โดยใช้แก้ปัญหาเกี่ยวกับข้อมูลที่มีลักษณะจำแนกกลุ่ม (Categorical Data) โดยใช้วิธีการที่เรียกว่า Categorical Method Variables (CMV) ซึ่งมีจุดมุ่งหมายในการจัดการกับข้อมูลจำแนกกลุ่มในโมเดลการวิเคราะห์

องค์ประกอบ ประมาณค่าโดยใช้เมตริกซ์สหสัมพันธ์ Polychoric และ Polyserial หรือ Biserial ระหว่างตัวแปรสังเกตได้ โดยใช้เมตริกซ์สหสัมพันธ์ Polychoric กับตัวแปรระดับเรียงอันดับ (Ordinal Scale) ใช้เมตริกซ์สหสัมพันธ์ Polyserial ระหว่างตัวแปรหนึ่งอยู่ในระดับเรียงอันดับและอีกตัวแปรหนึ่งเป็นข้อมูลต่อเนื่อง และเมตริกซ์สหสัมพันธ์ Biserial ระหว่างตัวแปรทวิ (Dichotomous) และอีกตัวแปรหนึ่งเป็นข้อมูลต่อเนื่อง ซึ่งในงานวิจัยนี้ผู้วิจัยจะกล่าวถึง เฉพาะเมตริกซ์สหสัมพันธ์ Polychoric กับข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordinal Categorical Data) โดยเฉพาะ Likert Items

แม้ว่าวิธีการ CMV จะมีประสิทธิภาพเมื่อข้อมูลมีลักษณะจำแนกกลุ่มที่มีการแจกแจง ไม่เป็น โกล่งปกติแต่วิธีการนี้ก็ยังมีความต้องการกลุ่มตัวอย่างขนาดใหญ่ ดังนั้น Muthén, du Toit and Spisic (1997 cited in Trierweiler, 2009, pp. 29 – 30) จึงได้พัฒนาวิธีการที่มีความแข็งแกร่งต่อขนาด กลุ่มตัวอย่างขนาดเล็ก เรียกว่า Robust Weighted Least Square ซึ่งแบ่งเป็น 2 วิธี คือ WLSM (Weighted Least Square with Mean Adjusted χ^2) และ WLSMV (Weighted Least Square with Mean – variance Adjusted χ^2) ซึ่งวิธีการ WLSM และ WLSMV มีความแตกต่างจาก WLS/ ADF ที่การใช้เมตริกซ์ความแปรปรวนร่วม และสมาชิกในเมตริกซ์ที่ใช้ถ่วงน้ำหนักในแนวทแยง (Brown, 2006 cited in Byrne, 2012, p. 132) รายละเอียดดังตารางที่ 2 – 4

ตารางที่ 2 – 4 เปรียบเทียบวิธีการประมาณค่าแบบ WLS, WLSM และ WLSMV (Finney & DiStefano, 2006, p. 293)

วิธีการ	คำอธิบาย	การประมาณ ค่า χ^2	การประมาณ ค่าพารามิเตอร์	ค่าเบี่ยงเบน มาตรฐาน	การนำไปใช้
WLS	1. ประมาณ	ใช้อินเวอร์ส	ใช้เมตริกซ์	ใช้อินเวอร์ส	ตัวแปรภายใน
	ค่าพารามิเตอร์ด้วย	ของเมตริกซ์	ถ่วงน้ำหนัก	ของเมตริกซ์	เป็นแบบ
	วิธีกำลังสองน้อย	ถ่วงน้ำหนัก	แบบเต็มรูป	ถ่วงน้ำหนัก	จำแนกกลุ่ม
	ที่สุดถ่วงน้ำหนัก	แบบเต็มรูป	(Full Weight	แบบเต็มรูป	(Categorical)
	2. ประมาณค่า	(Full Weight	Matrix Used)	(Full Weight	หรือต่อเนื่อง
	Chi – square และ	Matrix Used		Matrix Used	(Continuous)
	ค่าเบี่ยงเบน	and		and Inverted)	
	มาตรฐานแบบปกติ	Inverted)			

ตารางที่ 2 – 4 (ต่อ)

วิธีการ	คำอธิบาย	การประมาณ ค่า χ^2	การประมาณ ค่าพารามิเตอร์	ค่าเบี่ยงเบน มาตรฐาน	การนำไปใช้
WLSM	1. ประมาณ ค่าพารามิเตอร์ด้วย วิธีกำลังสองน้อย ที่สุดถ่วงน้ำหนัก	ใช้เมตริกซ์ ถ่วงน้ำหนัก แบบเต็มรูป (Full Weight	ใช้เมตริกซ์ ถ่วงน้ำหนัก แนวทแยง (Diagonal	ใช้เมตริกซ์ ถ่วงน้ำหนัก แบบเต็มรูป (Full Weight	ตัวแปรภายใน เป็นแบบ จำแนกกลุ่ม อย่างน้อย 1
	2. Mean Adjusted χ^2	Matrix but not Inverted)	Weight Matrix Used	Matrix but not Inverted)	กลุ่ม (At Least 1 Categorical)
	3. Scale Standard Error				
WLSMV	1. ประมาณ ค่าพารามิเตอร์ด้วย วิธีกำลังสองน้อย ที่สุดถ่วงน้ำหนัก	ใช้เมตริกซ์ ถ่วงน้ำหนัก แบบเต็มรูป (Full Weight	ใช้เมตริกซ์ ถ่วงน้ำหนัก แนวทแยง (Diagonal	ใช้เมตริกซ์ ถ่วงน้ำหนัก แบบเต็มรูป (Full Weight	ตัวแปรภายใน เป็นแบบ จำแนกกลุ่ม อย่างน้อย 1
	2. Mean and Variance – Adjusted χ^2)	Matrix but not Inverted)	Weight Matrix Used	Matrix but not Inverted)	กลุ่ม (At Least 1 Categorical)
	3. Scale Standard Error				

วิธีการประมาณค่าแบบ Robust Weight Least Square (WLSM และ WLSMV)

เป็นเทคนิคที่คำนวณในโปรแกรม Mplus มีงานวิจัยจำนวนจำกัดที่ศึกษาการประมาณค่าด้วยวิธี Robust WLS ซึ่งผลการศึกษาพบว่า การประมาณค่าแบบนี้ได้ผลดีกว่าวิธีการ WLS แบบธรรมดา แม้มีเงื่อนไขเกี่ยวกับโมเดลขนาดใหญ่ (15 ตัวแปรขึ้นไป) หรือกลุ่มตัวอย่างขนาดเล็ก จะให้ค่า Chi – square และค่าเบี่ยงเบนมาตรฐานที่มีความเอนเอียงน้อย (Muthén, 1993 cited in Finney & DiStefano, 2006, p. 294) วิธี WLSM จะให้ค่าความคลาดเคลื่อนประเภท 1 (Type I Error) สูงกว่าวิธี WLSMV และ วิธี WLSMV มีแนวโน้มให้ผลลัพธ์ที่ดีกว่าภายใต้เงื่อนไขกลุ่มตัวอย่างขนาดเล็ก ($n = 200$) และเมื่อตัวแปรมีการแจกแจงไม่เป็นโค้งปกติแบบเบ้ (Muthén et al., 1997 cited in Finney & DiStefano, 2006, p. 294)

Brown (2006 cited in Byrne, 2012, p. 132) กล่าวว่า วิธีการ WLSMV เป็นวิธีการประมาณค่าพารามิเตอร์ที่เหมาะสมที่สุดในการวิเคราะห์โมเดลองค์ประกอบเชิงยืนยันสำหรับข้อมูลจำแนกกลุ่ม

ข้อเสนอแนะสำหรับการจัดการข้อมูลที่มีการแจกแจงไม่เป็นโค้งปกติและข้อมูลที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ (Non – normality and Ordered Categorical Data) นำเสนอได้ดังตารางที่ 2 – 5

ตารางที่ 2 – 5 ข้อเสนอแนะสำหรับการจัดการข้อมูลที่มีการแจกแจงไม่เป็นโค้งปกติและข้อมูลที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ

ประเภทข้อมูล (Type of Data)	ข้อเสนอแนะ (Suggestion)
ข้อมูลแบบต่อเนื่อง (Continuous Data)	
1. มีการแจกแจงแบบปกติโดยประมาณ	ใช้การประมาณค่าพารามิเตอร์ด้วยวิธี ML
2. การแจกแจงไม่เป็นโค้งปกติ ในระดับเล็กน้อยถึงปานกลาง	1. ใช้วิธีการประมาณค่า ML 2. ใช้วิธี S – B Scaling (Robust ML) เพื่อประมาณค่า χ^2 และ SE
3. การแจกแจงไม่เป็นโค้งปกติ ในระดับมาก	1. ใช้วิธี S – B Scaling (Robust ML) 2. ใช้วิธี Bootstrapping
ข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data)	
1. มีการแจกแจงแบบปกติโดยประมาณ	1. ใช้วิธีการประมาณค่าแบบ WLSMV ใน Mplus
2. การแจกแจงไม่เป็นโค้งปกติ ในระดับเล็กน้อยถึงปานกลาง	2. ใช้วิธีการประมาณค่า ML ถ้าจำแนกกลุ่มอย่างน้อย 5 ระดับ (5 Categories) 3. ใช้วิธี S – B Scaling ถ้าจำแนกกลุ่มอย่างน้อย 4 ระดับ (4 Categories)
3. การแจกแจงไม่เป็นโค้งปกติ ในระดับมาก หรือ การจำแนกกลุ่มจำนวนน้อย (Very Few Categories) เช่น น้อยกว่า 5 ระดับ	1. ใช้วิธีการประมาณค่าแบบ WLSMV ใน Mplus 2. ใช้วิธี S – B Scaling ในกรณีที่ Mplus ไม่สามารถประมาณค่าได้

3. วิธี Bayesian

Muthén and Asparouhov (2011, pp. 5 – 6) กล่าวว่า วิธีการประมาณค่าแบบ Bayesian ในเทคนิค CFA เริ่มได้รับความนิยมมากขึ้น ซึ่ง Muthén and Muthén (2010 cited in Muthén & Asparouhov, 2011, p. 5) ได้พัฒนาเทคนิคการประมาณค่าลงในโปรแกรม Mplus ซึ่งมีการวิเคราะห์ที่สะดวก ง่าย และความน่าสนใจของวิธีการประมาณค่าแบบ Bayesian มีดังนี้ คือ

1. สามารถศึกษาค่าประมาณของพารามิเตอร์ได้มากขึ้น เช่น เมื่อข้อมูลมีการแจกแจงไม่เป็นโค้งปกติ วิธีการ ML จะประมาณค่าพารามิเตอร์และค่าความคลาดเคลื่อนมาตรฐาน โดยให้ข้อสมมติ (Assume) ว่าข้อมูลมีการแจกแจงเป็นโค้งปกติบนฐานของกลุ่มตัวอย่างขนาดใหญ่ และมีช่วงเชื่อมั่น (Confidential Interval) เท่ากับ $\text{Estimate} \pm 1.96 \times SE$ ในขณะที่วิธีการ Bayesian ไม่มีข้อจำกัดเรื่องกลุ่มตัวอย่างขนาดใหญ่และสามารถให้ค่าประมาณได้ในทุกรูปแบบการแจกแจง โดยอ้างอิงการแจกแจงภายหลัง (Posterior Distribution) และช่วงเชื่อมั่นขึ้นอยู่กับค่าเปอร์เซ็นต์ไทล์ของการแจกแจงภายหลัง โดยค่าพารามิเตอร์ที่ได้และสถิติทดสอบมีความแกร่งต่อการแจกแจงที่ไม่เป็นโค้งปกติ
2. มีความแกร่งต่อกลุ่มตัวอย่างขนาดเล็ก สามารถประมาณค่าได้ดีโดยไม่มีข้อจำกัดเรื่องขนาดกลุ่มตัวอย่าง
3. ระยะเวลาในการคำนวณโดยคอมพิวเตอร์มีความรวดเร็วกว่าวิธีการอื่น ๆ โดยเฉพาะเมื่อเทียบกับวิธีการ ML จากการศึกษาของ Browne and Draper (2006, p. 505 cited in Muthén & Asparouhov, 2011, p. 6) พบว่า ผลลัพธ์จากการประมาณค่าด้วยวิธีการ ML และวิธี Bayesian ให้ผลลัพธ์ที่ไม่แตกต่างกัน แต่วิธีการ Bayesian สามารถคำนวณด้วยระยะเวลาที่รวดเร็วกว่าสอดคล้องกับการศึกษาของ Asparouhov and Muthén (2010, p. 28) ในการเปรียบเทียบการประมาณค่าในโมเดล SEM โดยใช้การคำนวณด้วยโปรแกรม Mplus พบว่า วิธีการ Bayesian ใช้เวลาเฉลี่ยเร็วกว่าวิธีการ ML ประมาณ 200 เท่า
4. สามารถวิเคราะห์ข้อมูลสำหรับโมเดลแบบใหม่ ๆ ได้ดี เช่น โมเดลที่มีจำนวนพารามิเตอร์มาก ๆ

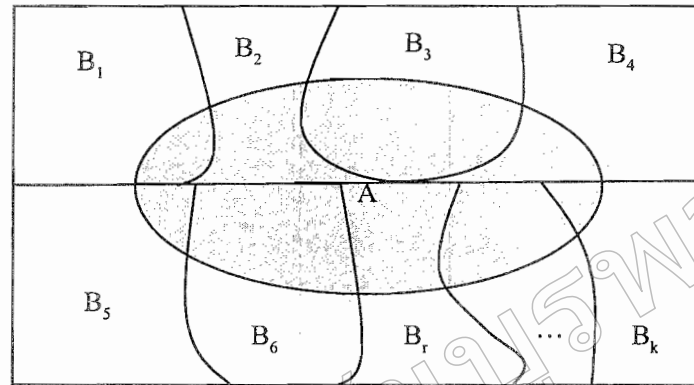
ทฤษฎีบทของเบย์ (Bayes' Theorem)

ทฤษฎีบทของเบย์เป็นทฤษฎีสำหรับหาค่าความน่าจะเป็นแบบมีเงื่อนไขของเหตุการณ์ที่สนใจจากเหตุการณ์ที่เป็นเงื่อนไข

เมื่อกำหนดให้ $B_1, B_2, \dots, B_k$ เป็นเหตุการณ์ที่ไม่เกิดขึ้นร่วมกัน (Mutually Exclusive Event) และครอบคลุมทั่ว Sample Space (Exhaustive) เมื่อกำหนดให้ A เป็นเหตุการณ์ใด ๆ ใน Sample Space

$$B_i \cap B_j = \emptyset \text{ เมื่อ } i \neq j$$

$$B_1 \cup B_2 \cup \dots \cup B_k = \text{Sample Space}$$



ภาพที่ 2-3 เหตุการณ์ที่สนใจจากเหตุการณ์ที่เป็นเงื่อนไข (Freund, 1992, pp. 70-72)

จะได้ความน่าจะเป็นแบบมีเงื่อนไขของทฤษฎีของเบย์ (Freund, 1992, pp. 70-72; Muthén & Asparouhov, 2011, p. 8) ดังนี้

$$P(B_r|A) = \frac{P(B_r)P(A|B_r)}{\sum_{i=1}^k P(B_i)P(A|B_i)}, \text{ เมื่อ } r = 1, 2, \dots, k \text{ หรือ } P(B_r|A) = \frac{P(B_r \cap A)}{P(A)}$$

โดยที่

$P(A)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ A

$P(B_r)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ B_r

$P(B_i)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ B_i

$P(B_r|A)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ B_r เมื่อมีเหตุการณ์ A เกิดก่อนแล้ว

$P(A|B_r)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ A เมื่อมีเหตุการณ์ B_r เกิดก่อนแล้ว

$P(A|B_i)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ A เมื่อมีเหตุการณ์ B_i เกิดก่อนแล้ว

$P(B_r \cap A)$ คือ ความน่าจะเป็นที่จะเกิดเหตุการณ์ A และ B_r ขึ้นร่วมกัน

ฟังก์ชันมาร์จินัล (Marginal Function)

จากฟังก์ชันการแจกแจงร่วมกันระหว่าง 2 ตัวแปร ในบางครั้งเราอาจต้องการศึกษาเฉพาะฟังก์ชันการแจกแจงความน่าจะเป็นของตัวแปรสุ่มเพียง 1 ตัวเท่านั้น จึงจำเป็นต้องศึกษาการแจกแจงมาร์จินัลของตัวแปรสุ่มนั้น

ถ้า X และ Y เป็นตัวแปรสุ่ม 2 มิติ ที่มีฟังก์ชันความน่าจะเป็น $f(x, y)$ แล้ว การแจกแจงความน่าจะเป็นมาร์จินัลของตัวแปรสุ่ม X และ Y กำหนดโดยฟังก์ชัน $g(x)$ และ $h(y)$ ตามลำดับพิจารณา

$$g(x) = \int_{-\infty}^{\infty} f(x, y) dy$$

$$h(y) = \int_{-\infty}^{\infty} f(x, y) dx$$

$$f(x, y) = f(x|y)f(y)$$

$$f(x, y) = f(y|x)f(x)$$

$$f(x) = \int_{-\infty}^{\infty} f(x, y) dy$$

$$= \int_{-\infty}^{\infty} f(x|y)f(y) dy$$

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

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

โดยที่ $\int_{-\infty}^{\infty} g(x) dx = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x, y) dy dx = 1$

$$\int_{-\infty}^{\infty} h(y) dy = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x, y) dx dy = 1$$

$$P(a < X < b) = \int_a^b g(x) dx$$

$$= \int_a^b \int_{-\infty}^{\infty} f(x, y) dy dx$$

$$P(c < Y < d) = \int_c^d h(y) dy$$

$$= \int_c^d \int_{-\infty}^{\infty} f(x, y) dx dy$$

ฟังก์ชันไลค์ลิฮูด (Likelihood Function)

ถ้าให้ $X_1, X_2, \dots, X_n$ เป็นตัวอย่างสุ่มขนาด n จะเรียกฟังก์ชันความน่าจะเป็นร่วม (Joint Density Function) ของ $X_1, X_2, \dots, X_n$ ซึ่งคือ $f(x_1, x_2, \dots, x_n)$ เรียกว่า ฟังก์ชันไลค์ลิฮูด (Likelihood Function) ซึ่งเป็นวิธีการหาตัวประมาณค่าพารามิเตอร์ โดยตัวประมาณค่าที่ได้ มีคุณสมบัติหลายประการแก่การเป็นตัวประมาณค่า แทนด้วยสัญลักษณ์ว่า $L(\theta, x_1, x_2, \dots, x_n)$ หรือ $L(\theta)$

กล่าวคือ ถ้า $X_1, X_2, \dots, X_n$ เป็นตัวอย่างสุ่มขนาด n จากการแจกแจง $f(x, \theta)$ จะได้ฟังก์ชันไลค์ลิฮูดของตัวแปรสุ่มทั้ง n ตัวแปรนี้เป็น

$$\begin{aligned} L(\theta, x_1, x_2, \dots, x_n) &= L(\theta) \\ &= f(x_1, x_2, \dots, x_n, \theta) \\ &= f(x_1, \theta) f(x_2, \theta) \dots f(x_n, \theta) \text{ เมื่อ } x_1, x_2, \dots, x_n \text{ เป็นอิสระกัน} \\ &= \prod_{i=1}^n f(x_i, \theta) \end{aligned}$$

การแจกแจงก่อน (Prior Distribution)

$P(\theta)$ คือ การแจกแจงความน่าจะเป็นของพารามิเตอร์ θ อาจคิดว่าความน่าจะเป็นก่อนที่จะทำการเก็บรวบรวมข้อมูล โดยอาศัยความรู้เกี่ยวกับสิ่งที่ต้องการศึกษา ซึ่งได้มาจากการประสบการณ์การศึกษาหรือทำการวิจัย ก่อนที่จะทำการเก็บรวบรวมข้อมูล เรียกความรู้ที่ว่า การแจกแจงก่อน

พิจารณาฟังก์ชันเงื่อนไข

$$f(y|x) = \frac{f(x, y)}{f(x)}$$

$$f(x, y) = f(y|x) f(x)$$

$$f(x|y) = \frac{f(x, y)}{f(y)}$$

$$f(x, y) = f(x|y) f(y)$$

ฟังก์ชันมาร์จินัล

$$f(x, y) = f(x|y) f(y)$$

$$f(x, y) = f(y|x) f(x)$$

$$f(x) = \int f(x, y) dy$$

$$= \int f(x|y) f(y) dy$$

$$f(y) = \int f(x, y) dx$$

$$= \int f(y|x) f(x) dx$$

$P(X)$ เป็นการแจกแจงความน่าจะเป็นไม่จำกัดของข้อมูล ดังนั้น การแจกแจงมาร์จินัลของข้อมูล คือ $P(X) = \int P(X|\theta) P(\theta) d\theta$ ซึ่งก็คือ Prior Distribution

การแจกแจงภายหลัง (Posterior Distribution)

คือ ความน่าจะเป็นแบบมีเงื่อนไขเป็นการนำเอา $P(\theta)$ (Prior distribution) มาผนวกกับข้อมูลที่เก็บรวบรวมมาจะทำให้ได้พารามิเตอร์แบบใหม่ที่เรียกว่า Posterior Distribution หรือ $P(\theta|X)$ ซึ่งเป็นค่าความน่าจะเป็นที่ผสมผสานระหว่างความรู้ที่มีอยู่ในเรื่องนั้นกับข้อมูลที่เก็บรวบรวมได้ จะได้

$$\text{Posterior distribution } (P(\theta|X)) = \frac{P(\theta) \times P(x|\theta)}{\int P(\theta) \times P(x|\theta) d\theta}$$

ซึ่งเมื่อทำการอินทิเกรตค่า Normalizing Constant แล้วได้ค่าคงที่ ดังนั้น

$$\text{Posterior distribution } (P(\theta|X)) = P(\theta)P(X|\theta)$$

เมื่อมีค่าสังเกต จะได้ฟังก์ชันการแจกแจงภายหลังคือ

$$P(\theta|x_1) \propto P(\theta)P(x_1|\theta)$$

ทฤษฎีความน่าจะเป็นแบบมีเงื่อนไขของทฤษฎีของเบย์สามารถนำไปประยุกต์ใช้ในการแจกแจงของตัวแปรได้

ถ้าให้ $B_1, B_2, \dots, B_k$ เป็นเหตุการณ์ที่ถูกสังเกต n เหตุการณ์ โดยที่เหตุการณ์นี้ไม่เกิดร่วมกัน ซึ่งมีการแจกแจงความน่าจะเป็นคือ $P(B_i|A)$ ที่ขึ้นอยู่กับพารามิเตอร์ A และมีความน่าจะเป็นคือ $P(A)$ แล้วจะได้ความน่าจะเป็นแบบมีเงื่อนไขของ A เมื่อ $P(A) > 0$ คือ

$$P(B_i|A) = \frac{P(A|B_i)P(B_i)}{P(A)}, \text{ เมื่อ } P(A) = \sum_{i=1}^k P(B_i)P(A|B_i)$$

จากสมการดังกล่าว จะได้ว่า $P(B_i|A)$ คือ การแจกแจงภายหลัง (Posterior Distribution) ของเหตุการณ์ B_i และ $P(B_i)$ คือ การแจกแจงก่อน (Prior Distribution) ของเหตุการณ์ B_i ก่อนที่ข้อมูลจะถูกสังเกตหรือทดลอง และ $P(A|B_i)$ คือ Likelihood Function

กล่าวโดยสรุป แนวคิดการประมาณค่าแบบ Bayesian นี้ จะเป็นการอาศัยข้อมูลที่ได้จากการสังเกตมาเปลี่ยนแปลงความเชื่อดั้งเดิมเกี่ยวกับค่าที่แตกต่างกันของ P โดยการเปลี่ยนฟังก์ชันความน่าจะเป็นเบื้องต้นของ P ไปสู่ฟังก์ชันความน่าจะเป็นภายหลังของ P (Muthén & Asparouhov, 2011)

$$\text{Posterior Probability} = \frac{\text{Prior Probability} \times \text{Likelihood}}{\text{Marginal Distribution of Data}}$$

การประมาณค่าด้วยวิธี Bayesian (Bayesian Estimation)

กำหนดให้ $x_1, x_2, \dots, x_n$ เป็นตัวอย่างสุ่ม (Random Sample) จากประชากรที่มีฟังก์ชันความหนาแน่น (Probability Density Function) คือ $f(x|\theta)$ โดยที่ θ เป็นพารามิเตอร์ที่ไม่ทราบค่า และต้องการประมาณค่า

โดยทั่วไปแล้ว ในการประมาณค่าพารามิเตอร์ θ ของการแจกแจงของประชากร มักจะถือหลักว่าพารามิเตอร์ θ เป็นค่าคงที่แต่ไม่ทราบค่า และการประมาณค่า θ จะทำโดยใช้ตัวอย่างสุ่มจากการแจกแจงของประชากรนั้นๆ แต่ในบางครั้ง ผู้ที่ประมาณค่าอาจจะทราบข้อเท็จจริงบางอย่างเกี่ยวกับ θ ก่อนที่จะสุ่มตัวอย่าง ซึ่งหากนำข้อเท็จจริงนั้นมาใช้ให้เกิดประโยชน์ก็จะช่วยให้การประมาณค่านั้นให้ผลดียิ่งขึ้น

ในการประมาณค่าด้วยวิธี Bayesian จะถือว่าค่าพารามิเตอร์ θ เป็นค่าของตัวแปรสุ่ม Θ ที่มีการแจกแจงความน่าจะเป็นที่ทราบล่วงหน้า โดยแสดงได้ด้วยฟังก์ชันความหนาแน่น $g(\theta)$ และเรียกฟังก์ชัน $g(\theta)$ ว่าฟังก์ชันความหนาแน่นเบื้องต้น (Prior Probability Density Function) ของ Θ ดังนั้น ฟังก์ชัน $f(x|\theta)$ จึงเป็นฟังก์ชันความหนาแน่นอย่างมีเงื่อนไข (Condition Probability Density Function) ของ x เมื่อกำหนดค่า $\Theta = \theta$ ซึ่งเขียนแทนด้วย $f(x|\theta)$ และฟังก์ชันความหนาแน่นอย่างมีเงื่อนไขของ Θ เมื่อกำหนดค่า $X_1 = x_1, X_2 = x_2, \dots, X_n = x_n$ เรียกว่า ฟังก์ชันความหนาแน่นภายหลัง (Posterior Probability Density Function) ของ Θ ซึ่งกำหนดสัญลักษณ์ฟังก์ชัน $h(\theta | x_1, x_2, \dots, x_n)$ เป็นฟังก์ชันความหนาแน่นภายหลังของ Θ

จากกฎของเบย์ (Bays' Theorem)

$$\begin{aligned} P(B_i|A) &= P(A \cap B_i) / P(A) \\ &= P(A|B_i) P(B_i) / P(A), i = 1, 2, \dots, m \end{aligned}$$

$$\text{เมื่อ } P(B_i \cap B_j) = 0; i \neq j \text{ และ } P(B_1 \cup B_2 \cup \dots \cup B_m) = 1$$

จะได้ว่า

เมื่อ Θ เป็นตัวแปรสุ่มแบบต่อเนื่อง (Continuous Random Variable)

$$h(\theta | x_1, x_2, \dots, x_n) = \frac{\prod_{i=1}^n f(x_i | \theta) g(\theta)}{\int \prod_{i=1}^n f(x_i | \theta) g(\theta) d\theta}$$

เมื่อ θ เป็นตัวแปรสุ่มแบบไม่ต่อเนื่อง (Discrete Random Variable)

$$h(\theta | x_1, x_2, \dots, x_n) = \frac{\prod_{i=1}^n f(x_i | \theta) g(\theta)}{\sum_{\text{all } \theta} \prod_{i=1}^n f(x_i | \theta) g(\theta)}$$

นอกจากการประมาณ θ แล้ว เราอาจประมาณค่าของฟังก์ชันของ θ ได้อีกด้วย สมมติให้ $\tau(\theta)$ แทนฟังก์ชันของ θ ที่ต้องการประมาณค่า และให้ T เป็นฟังก์ชันของตัวอย่างสุ่ม $X_1, X_2, \dots, X_n$ โดย T จะใช้เป็นตัวประมาณของ $\tau(\theta)$ ถ้า $k(t | \theta)$ เป็นฟังก์ชันความหนาแน่นอย่างมีเงื่อนไขของ T เมื่อกำหนดค่า $\theta = \theta$ ได้ว่า

$$\begin{aligned} h(\theta | t) &= h(\theta, t) / h_1(t) \\ &= k(t | \theta) g(\theta) / h_1(t) \end{aligned}$$

เมื่อ $h(\theta, t)$ แทน ฟังก์ชันความหนาแน่นร่วม (Joint Probability Density Function)

ของ θ และ T

$h_1(t)$ แทน ฟังก์ชันความหนาแน่นมาร์จินัล (Marginal Probability Density Function)

ของ T

การประมาณค่าแบบ Bayesian เป็นการเลือกฟังก์ชันการตัดสินใจ (Decision Function: $d(t)$) เพื่อใช้ประมาณค่า $\tau(\theta)$ เมื่อทราบค่า t จากตัวอย่างสุ่ม โดยฟังก์ชันความหนาแน่น Posterior $h(\theta | t)$ และการเลือกฟังก์ชันการตัดสินใจจะพิจารณาโดยใช้ฟังก์ชันการสูญเสีย (Loss Function) เป็นเกณฑ์

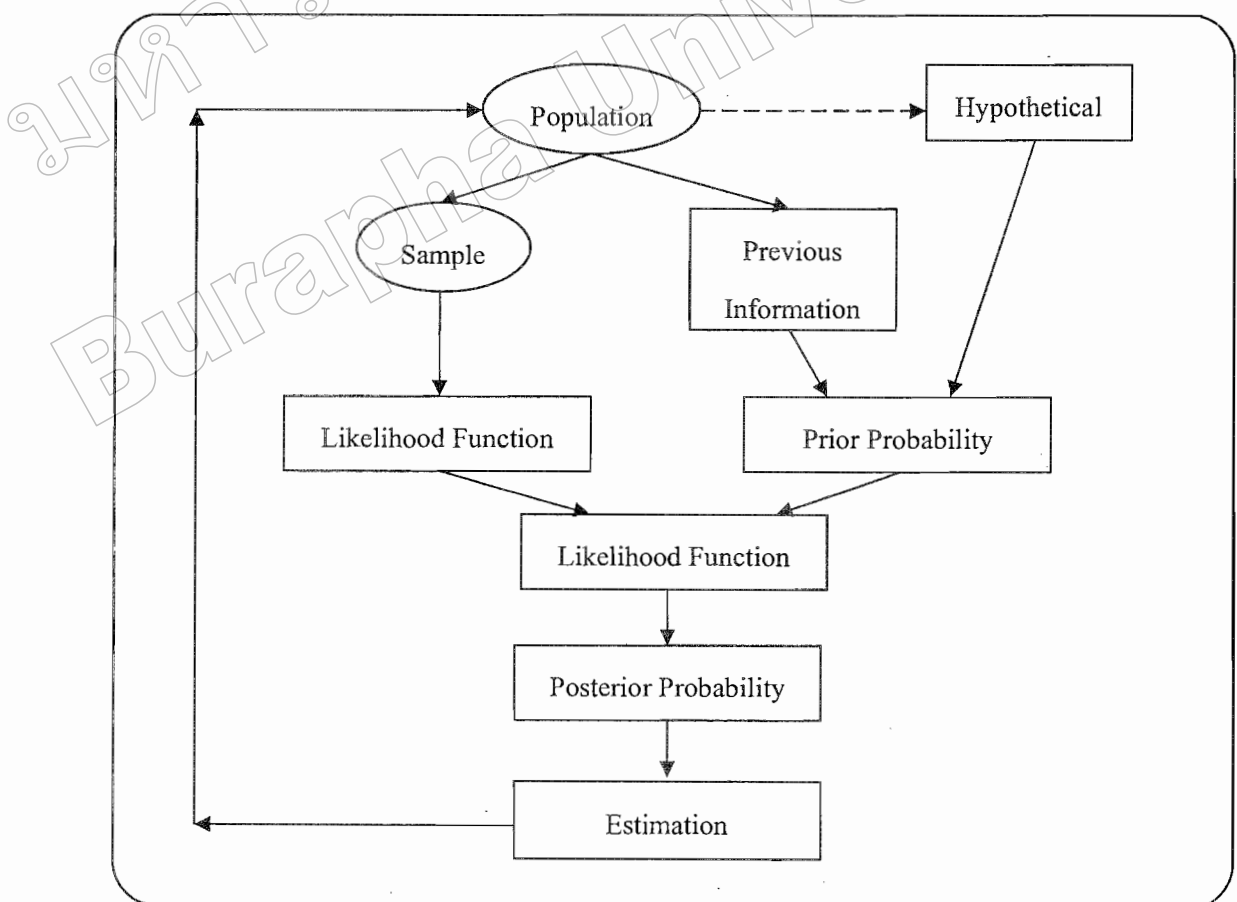
นิยาม ฟังก์ชันการสูญเสีย $L[\theta; d(t)]$ ได้แก่ ฟังก์ชันค่าจริงที่มีคุณสมบัติ ดังนี้

1. $L[\theta; d(t)] \geq 0$ ทุก $d(t)$ และ $\theta \in \Theta$
2. $L[\theta; d(t)] = 0$ เมื่อ $d(t) = \tau(\theta)$, โดยที่ $d(t)$ เป็นตัวประมาณของ $\tau(\theta)$

แนวทางหนึ่งในการเลือกใช้ฟังก์ชันการตัดสินใจ โดยขึ้นอยู่กับฟังก์ชันการสูญเสีย นั่นคือ การเลือกฟังก์ชันการตัดสินใจที่ทำให้ค่าคาดหวังของการสูญเสียอย่างมีเงื่อนไข (Conditional Expectation of the Loss) มีค่าต่ำที่สุด เมื่อกำหนด $T = t$ ตัวอย่างเช่น $L[\theta; d(t)] = [\theta - d(t)]^2$ ซึ่งเรียกว่า ฟังก์ชันการสูญเสียความคลาดเคลื่อนกำลังสอง (Squared Error Loss Function) จะได้ว่าตัวประมาณค่าแบบ Bayesian ของ θ คือ ค่าเฉลี่ยของการแจกแจงภายหลัง (Posterior Distribution) ของ θ หรือ $E(\theta | t)$ เนื่องจาก $E[\theta - d(t) | T = t]$ มีค่าต่ำที่สุดเมื่อ $d(t) = E(\theta | t)$

วิธีการและการอ้างอิงสถิติแบบ Bayesian (Bayesian Process and Inference)

ในกระบวนการอ้างอิงทางสถิติแบบ Bayesian จะถือว่าค่าพารามิเตอร์ของประชากรนั้นเป็นตัวแปร ซึ่งสถิติแบบดั้งเดิมถือว่าค่าพารามิเตอร์เป็นค่าคงที่ ดังนั้นในสถิติแบบ Bayesian ค่าพารามิเตอร์จะมีการแจกแจงก่อน (Prior Distribution) เมื่อทำการศึกษาจะทำการสุ่มตัวอย่างมาเพื่อศึกษาคุณลักษณะของประชากรที่ต้องการ อาศัยทฤษฎีเบย์รวมทั้งข้อมูลเบื้องต้นและข้อมูลที่ได้จากตัวอย่างเข้าด้วยกันจะได้การแจกแจงที่เรียกว่า การแจกแจงภายหลัง (Posterior Distribution) แล้วทำการประมาณค่าพารามิเตอร์จากการแจกแจงดังกล่าว ดังภาพที่ 2 – 4



ภาพที่ 2 – 4 การอ้างอิงทางสถิติแบบ Bayesian

ดัชนีที่ใช้วัดความสอดคล้องกลมกลืนของโมเดลสำหรับวิธีการประมาณค่าแบบ Bayesian คือ Posterior Predictive p – value ซึ่งเป็นดัชนีที่วัดความกลมกลืนของโมเดลไม่ใช่ค่าทดสอบนัยสำคัญทางสถิติ ถ้าโมเดลมีความสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ดีค่าที่ได้ควรมีค่าเข้าใกล้ 0 หรือน้อยกว่า .05 (Muthén & Asparouhov, 2011)

ประสิทธิภาพของการประมาณค่าพารามิเตอร์ (Efficiency of Parameter Estimates)

ความเอนเอียงสัมพัทธ์ของค่าประมาณพารามิเตอร์ (Parameter Estimates Relative Bias)

ความเอนเอียงสัมพัทธ์ (Relative Bias: RB) หรือ ค่าร้อยละของความเอนเอียง (Percentage Bias) ของค่าประมาณพารามิเตอร์ คือ ค่าร้อยละของความแตกต่างระหว่างค่าประมาณพารามิเตอร์ที่ได้จากกลุ่มตัวอย่างและค่าจริงของพารามิเตอร์จากโมเดลประชากร (Bandalos, 2006, p.401; Flora & Curran, 2004, p. 473; Trierweiler, 2009, p. 75; Chumney, 2012, p. 31) ซึ่งสามารถคำนวณเป็นค่าร้อยละของความเอนเอียงสัมพัทธ์ได้จากสูตร

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

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

θ_i คือ ค่าพารามิเตอร์จากประชากร

การคำนวณความเอนเอียงสัมพัทธ์ของค่าประมาณของพารามิเตอร์ที่เป็นค่ามาตรฐานในโมเดล CFA ในการวิจัยครั้งนี้ ได้แก่ ค่าน้ำหนักองค์ประกอบ (Factor Loading) ค่าปฏิสัมพันธ์ขององค์ประกอบ (Factor Correlation) จะคำนวณโดยการหาค่าเฉลี่ยจากการทดลองซ้ำ (Replication) ทั้งหมดที่ทำการศึกษา ในบางครั้งหากผลต่างค่าประมาณพารามิเตอร์จากกลุ่มตัวอย่างและจากประชากรมีค่าเป็นลบ หรือเป็นบวกที่ต่างกันมาก ๆ แต่เมื่อนำมาหาค่าเฉลี่ยแล้ว พบว่ามีค่าความเอนเอียงเป็นศูนย์ หรือไม่มีความเอนเอียง การแปลความหมายอาจผิดพลาดได้ ดังนั้นหากผู้วิจัยไม่สนใจทิศทางของความเอนเอียง ควรใส่ค่าสัมบูรณ์สำหรับตัวเศษในสูตร คือ $|\hat{\theta}_{ij} - \theta_i|$

เกณฑ์การพิจารณาความเอนเอียงสัมพัทธ์ของค่าประมาณพารามิเตอร์ มีเกณฑ์ดังต่อไปนี้ (Curran, West, & Finch, 1996; Bandalos, 2006, p.402; Flora & Curran, 2004, p. 473; Kaplan, 1989 cited in Trierweiler, 2009, p. 75)

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

ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias)

ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน คำนวณจากการหาค่าเฉลี่ยของความแตกต่างของค่าประมาณความคลาดเคลื่อนมาตรฐานจากกลุ่มตัวอย่างและค่าความคลาดเคลื่อนมาตรฐานจากโมเดลประชากร จากจำนวนชุดข้อมูล (Replication) ทั้งหมด ที่ศึกษาเช่นเดียวกับค่าความเอนเอียงของค่าประมาณพารามิเตอร์ โดยคำนวณได้จากสูตร

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

$SE(\hat{\theta}_i)$ คือ ค่าประมาณของความคลาดเคลื่อนมาตรฐานของพารามิเตอร์จากประชากร

$SE(\hat{\theta}_i)_j$ คือ ค่าประมาณของความคลาดเคลื่อนมาตรฐานจากกลุ่มตัวอย่าง สำหรับจำนวนชุดข้อมูล (Replication) ที่ j

เกณฑ์สำหรับค่าความเอนเอียงสัมพัทธ์ (Relative Bias) ของของความคลาดเคลื่อนมาตรฐาน (Curran, West, & Finch, 1996; Bandalos, 2006, p.403; Flora & Curran, 2004, p. 473; Kaplan, 1989 cited in Trierweiler, 2009, p. 77) คือ

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

การประเมินความสอดคล้องของโมเดล (Assessment of Model Fit)

การตรวจสอบโมเดลการวัด (Measurement Model) (Hair et al., 2010, pp. 677 – 680)

โมเดลการวัด เป็นโมเดลที่ใช้ตัวแปรสังเกตได้วัดตัวแปรแฝง ดังนั้นในการแปลผลการวิเคราะห์ควรพิจารณาด้วยว่าตัวแปรสังเกตได้วัดตัวแปรแฝงได้มากน้อยเพียงใด การพิจารณาประสิทธิภาพของโมเดลการวัดจะต้องพิจารณา 3 ลักษณะ คือ

1. ความเป็นเอกมิติ (Unidimensionality) ของชุดตัวแปรสังเกตได้ หมายถึง การที่ชุดตัววัดหรือชุดข้อคำถามที่ใช้วัดตัวแปรแฝงสามารถรวมตัวกันได้องค์ประกอบเดียว หรือ

มีองค์ประกอบที่เด่นที่สุด (Dominant) เพียงองค์ประกอบเดียว การตรวจสอบความเป็นเอกมิติ จะต้องใช้การวิเคราะห์องค์ประกอบเชิงยืนยัน (CFA) แสดงให้เห็นว่าชุดตัวชี้วัดมีคุณลักษณะแฝงเดียวที่ต้องการวัด ซึ่งเป็นส่วนสำคัญที่แสดงถึงความตรงเชิงจำแนก (Discriminant Validity) หมายถึง ขอบเขตที่ตัวแปรแฝงตัวหนึ่งแตกต่างจากตัวแปรแฝงตัวอื่น ๆ การตรวจสอบความตรงเชิงจำแนกยังสามารถตรวจสอบได้ด้วยค่า Average Variance Extracted ที่มากกว่าค่ากำลังสองของสัมประสิทธิ์สหสัมพันธ์ (Square Correlation) ระหว่างตัวแปรแฝง

2. ความเชื่อมั่น (Reliability) ของชุดตัวแปรสังเกตได้ หมายถึง ความคงเส้นคงวาหรือความสอดคล้องภายใน (Internal Consistency) ของชุดตัวแปรสังเกตได้ ในแต่ละกลุ่มตัวแปรแฝงเป็นค่าที่แสดงถึงระดับความเป็นตัวแทนตัวแปรแฝงของชุดตัวแปรสังเกตได้ หรือความแม่นยำของโมเดลการวัด แบ่งเป็น

2.1 ความเชื่อมั่นของข้อคำถาม (Item Reliability) เป็นค่าความเชื่อมั่นที่แสดงให้เห็นคุณภาพของข้อคำถามรายข้อ ในผลการวิเคราะห์ค่าความเชื่อมั่นของข้อคำถามได้จากค่า Square Multiple Correlations ซึ่งเป็นสัดส่วนความแปรปรวนของตัวแปรที่อธิบายได้โดยตัวแปรแฝงซึ่งมีค่าเท่ากับค่าการร่วมกัน (Communality) ในการวิเคราะห์องค์ประกอบเชิงสำรวจ

2.2 ความเชื่อมั่นของตัวแปรแฝง (Construct Reliability: ρ_c) แสดงถึงความตรงในการรวมตัว (Convergent Validity) ซึ่งหมายถึง สัดส่วนความแปรปรวนร่วมกันของตัวแปรสังเกตได้ทั้งหมดในตัวแปรแฝงเดียวกันคำนวณอีกครั้งหนึ่ง จากสูตร

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

เมื่อ λ_i คือ น้ำหนักองค์ประกอบมาตรฐาน

δ_i คือ ความแปรปรวนของความคลาดเคลื่อนมาตรฐาน

การคำนวณควรแยกแต่ละตัวแปรแฝงและค่า ρ_c ที่ได้ควรมากกว่า .60 (Hair et al., 2010, p. 680)

2.3 ค่าเฉลี่ยของความแปรปรวนที่ถูกสกัดได้ (Average Variance Extracted: ρ_v) เป็นค่าเฉลี่ยของความแปรปรวนของตัวแปรสังเกตที่อธิบายได้ด้วยตัวแปรแฝงเมื่อเทียบกับความแปรปรวนของความคลาดเคลื่อนในการวัด ซึ่งมีค่าเทียบเท่ากับค่าไอเก้น (Eigen Value) ในการวิเคราะห์องค์ประกอบเชิงสำรวจ และคำนวณจากสูตร

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

เมื่อ λ_i คือ น้ำหนักองค์ประกอบมาตรฐาน

δ_i คือ ความแปรปรวนของความคลาดเคลื่อนมาตรฐาน

การคำนวณควรแยกแต่ละตัวแปรแฝง และค่า ρ_v ที่ได้ควรมากกว่า .50 (Hair et al., 2010, p. 680)

3. ความตรง (Validity) หมายถึง ความสามารถของตัวแปรสังเกตได้ที่ใช้วัดตัวแปรแฝงในโมเดล โดยพิจารณาจากความมีนัยสำคัญของน้ำหนักองค์ประกอบ (Factor Loading: λ) และค่าคลาดเคลื่อนมาตรฐาน (Standard Error) ซึ่งค่าน้ำหนักองค์ประกอบควรมีค่าสูง คือ มีค่ามาตรฐานตั้ง .70 ขึ้นไป และมีนัยสำคัญทางสถิติ (Hair et al., 2010, p. 679) ถ้าค่าน้ำหนักองค์ประกอบไม่มีนัยสำคัญแสดงว่าค่าคลาดเคลื่อนมาตรฐานมีขนาดใหญ่และโมเดลการวิจัยยังไม่ดีพอ ค่าสถิติ t (t-value) มีค่ามากกว่า 1.96 (ที่ระดับนัยสำคัญทางสถิติ .05) หรือมีค่ามากกว่า 2.38 (ที่ระดับนัยสำคัญทางสถิติ .01) นอกจากนี้ยังสามารถเปรียบเทียบความสำคัญของตัวแปรสังเกตได้ว่าตัวแปรใดใช้วัดตัวแปรแฝงได้ดีที่สุด โดยการเปรียบเทียบน้ำหนักองค์ประกอบมาตรฐาน (Standardized Loading) ตัวแปรสังเกตได้ที่มีความสำคัญมาก จะมีค่าน้ำหนักองค์ประกอบมาตรฐานสูง

การตรวจสอบโมเดลโครงสร้าง (Structural Model)

การตรวจสอบโมเดลโครงสร้าง เป็นการตรวจสอบความสัมพันธ์ระหว่างโมเดลที่พัฒนาขึ้นตามทฤษฎีกับโมเดลจากข้อมูลเชิงประจักษ์ โดยแนวคิดของโมเดลสมการโครงสร้างเป็นการวิเคราะห์เพื่อเปรียบเทียบความสัมพันธ์กันระหว่างสองสิ่ง สิ่งแรกเป็นความสัมพันธ์ทางสถิติระหว่างกลุ่มตัวแปรที่สามารถสังเกตและวัดได้จากข้อมูลที่รวบรวมมาจากกลุ่มตัวอย่าง สิ่งที่สองเป็นความสัมพันธ์ทางสถิติระหว่างกลุ่มตัวแปรที่ถูกประมาณค่าขึ้นมาโดยอาศัยความรู้ทางทฤษฎีและจากการทบทวนวรรณกรรมที่เกี่ยวข้อง ดังนั้นจึงเป็นการวิเคราะห์เปรียบเทียบเมตริกซ์ความแปรปรวนและความแปรปรวนร่วม 2 ชุด คือ

1. เมตริกซ์ความแปรปรวนและความแปรปรวนร่วมของกลุ่มตัวอย่าง (Sample Variance – Covariance Matrix) สัญลักษณ์ คือ S
2. เมตริกซ์ความแปรปรวนและความแปรปรวนร่วมของประชากร (Predicted Variance – Covariance Matrix) เป็นเมตริกซ์ที่ประมาณค่าขึ้นจากตัวแปรตามที่ได้จากการทำนายในโมเดล สัญลักษณ์ คือ Σ

การเปรียบเทียบระหว่างเมตริกซ์ความแปรปรวนและความแปรปรวนร่วมทั้งสองนี้ ทำให้โมเดลสมการ โครงสร้างมีชื่อเรียกอีกชื่อหนึ่งว่า โมเดลโครงสร้างความแปรปรวนร่วม (Covariance Structural Model)

เมตริกซ์ความแปรปรวนและความแปรปรวนร่วมของกลุ่มตัวอย่าง ก็คือ เมตริกซ์ความแปรปรวนและความแปรปรวนร่วมระหว่างกลุ่มตัวแปรสังเกตได้ที่วัดจากกลุ่มตัวอย่าง ซึ่งได้มาจากการคำนวณ โดยใช้ตัวแปรที่มีหน่วยการวัดตามความเป็นจริงในการเก็บรวบรวมข้อมูล กรณีนี้ตัวแปรต่าง ๆ อาจมีหน่วยการวัดที่แตกต่างกันได้

ในการวิเคราะห์เพื่อการประมาณค่าพารามิเตอร์ นักวิจัยสามารถใช้เมตริกซ์สหสัมพันธ์ระหว่างตัวแปรแทนเมตริกซ์ความแปรปรวนและความแปรปรวนร่วมได้ เมตริกซ์สหสัมพันธ์เป็นเมตริกซ์ที่ได้จากการทำให้สมาชิกในเมตริกซ์ความแปรปรวนและความแปรปรวนร่วม มีหน่วยเป็นมาตรฐานเดียวกันด้วยการหารด้วยส่วนเบี่ยงเบนมาตรฐาน สมาชิกในเมตริกซ์ทุกตัว จะมีค่าสูงสุดเท่ากับ ± 1 และค่าต่ำสุดเท่ากับ 0

ความแปรปรวนร่วมและสัมประสิทธิ์สหสัมพันธ์มีความสัมพันธ์กันทางสถิติโดยตรง ดังนี้

$$r_{xy} = \frac{S_{xy}}{\sqrt{S_x^2} \sqrt{S_y^2}} \quad \text{หรือ} \quad r_{xy} = \frac{S_{xy}}{S_x S_y}$$

เมื่อ S_{xy} คือ ความแปรปรวนร่วม

S_x^2 และ S_y^2 คือ ความแปรปรวนของตัวแปร x และตัวแปร y

ดัชนีที่ใช้ในการตรวจสอบความสอดคล้องกลมกลืนระหว่างโมเดลตามสมมติฐาน กับข้อมูลเชิงประจักษ์ (Fit Indices)

ค่าสถิติในกลุ่มนี้ใช้ในการตรวจสอบความตรงของโมเดลเป็นภาพรวมทั้งโมเดล ในการวิเคราะห์โมเดลสมการเชิงโครงสร้าง โปรแกรมจะประเมินความสอดคล้องกลมกลืนของโมเดลกับข้อมูลเชิงประจักษ์ แล้วรายงานค่าดัชนีต่าง ๆ ในรายงานผลการวิเคราะห์ (Print Out) ค่าดัชนีเหล่านั้นจะแสดงว่าโดยภาพรวมโมเดลสมการเชิงโครงสร้างสอดคล้องกลมกลืนของโมเดลกับข้อมูลเชิงประจักษ์เพียงใด ดัชนีที่ใช้บอกความสอดคล้องกลมกลืนของโมเดลมีหลายตัว แต่ไม่มีดัชนีตัวใดตัวหนึ่งที่ดีกว่าตัวอื่น ๆ เพราะค่าดัชนีต่าง ๆ แต่ละตัวใช้ในแต่ละกรณี เช่น ขนาดกลุ่มตัวอย่าง วิธีการประมาณค่า ความซับซ้อนของโมเดล การไม่เป็นไปตามข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงปกติพหุนาม (Multivariate Normal Distribution) จำนวนตัวแปร

สังเกตได้ หรือหลาย ๆ กรณีรวมกัน ดัชนีวัดความสอดคล้องกลมกลืนของโมเดล (Schumacker & Lomax, 2010, pp. 73 – 92) ประกอบด้วย

1. Absolute Fit Indices

1.1 การทดสอบแบบ Overall Fit ด้วย Chi – square Test (χ^2 – test)

การทดสอบด้วยค่าสถิติ Chi – square เป็นดัชนีที่ใช้แพร่หลายในการตรวจสอบความสอดคล้องกลมกลืนของโมเดลกับข้อมูลเชิงประจักษ์ โดยภาพรวมค่า Chi – square คำนวณจากผลคูณระหว่าง Minimum Fit Function (F_{min}) กับ $n - 1$ เมื่อ n แทน ขนาดกลุ่มตัวอย่าง มีองศาของความเป็นอิสระ (df) เท่ากับ $[k(k + 1)/ 2] - t$ เมื่อ k แทน จำนวนตัวแปรสังเกตได้ และ t แทนจำนวนพารามิเตอร์ในโมเดลที่ต้องการประมาณค่า (Kline, 2005, p. 135) สมมติฐานของการทดสอบ คือ $H_0: S = \hat{\Sigma}$ เมื่อ S แทน เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมของข้อมูลเชิงประจักษ์ และ $\hat{\Sigma}$ แทน เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วมของตัวแปรสังเกตได้ที่ประมาณค่ามาจากระเบียง ถ้าค่า Chi – square มีนัยสำคัญแสดงว่า โมเดลกับข้อมูลเชิงประจักษ์ไม่สอดคล้องกลมกลืนกัน ดังนั้นการทดสอบด้วย Chi – square จึงเป็นการพิจารณาถึงความสอดคล้องมากกว่าทดสอบว่าโมเดลที่กำหนดขึ้นมานั้นถูกต้องหรือไม่ การใช้ Chi – square เป็นค่าสถิติทดสอบวัดความสอดคล้องกลมกลืนต้องใช้ด้วยความระมัดระวัง เพราะค่าสถิติ Chi – square มีข้อด้อยเบื้องต้นอยู่ 4 ประการ (Byrne, 2012, p. 67) คือ 1) ตัวแปรสังเกตได้ต้องมีการแจกแจงเป็นโค้งปกติ เพราะถ้าตัวแปรสังเกตได้มีการแจกแจงแบบโด่งสูง (Leptokurtic) จะทำให้ค่า Chi – square สูงกว่าความเป็นจริงทำให้มีโอกาสปฏิเสธสมมติฐานศูนย์ได้มาก ส่วนข้อมูลที่มีการแจกแจงแบบแบนราบ (Platykurtic) ก็จะทำให้ Chi – square ต่ำกว่าความเป็นจริง ถ้าข้อมูลมีความเบ้สูงก็จะทำให้ Chi – square สูงกว่าปกติ (Broom, 1981, p. 81 cited in Bollen, 1989, pp. 266 – 267) 2) การวิเคราะห์ข้อมูลต้องใช้เมตริกซ์ความแปรปรวน – ความแปรปรวนร่วม 3) ขนาดกลุ่มตัวอย่างต้องมีขนาดที่เหมาะสม เพราะ Chi – square ขึ้นอยู่กับขนาดของกลุ่มตัวอย่าง เพราะกลุ่มตัวอย่างที่มีขนาดใหญ่มาก ๆ ก็จะทำให้ค่า Chi – square สูงมาก และถ้ากลุ่มตัวอย่างมีขนาดเล็ก (ต่ำกว่า 100) ก็จะทำให้ค่า Chi – square มีแนวโน้มต่ำกว่าความเป็นจริง จนอาจทำให้สรุปผลไม่ถูกต้องได้ (Schumacker & Lomax, 2010, p. 86; Kline, 2005, p. 136) ดังนั้นจึงแก้ไขด้วยการพิจารณาค่า Chi – square สัมพัทธ์ (χ^2 / df) ซึ่งควรมีค่าน้อยกว่า 2.00 หรือในบางตำรากล่าวว่าควรมีค่าน้อยกว่า 5.00 (Bollen, 1989, p. 278) และ 4) ฟังก์ชันความกลมกลืนมีค่าเป็นศูนย์จริงตามสมมติฐานที่ใช้ทดสอบ Chi – square นั้น Schumacker and Lomax (2010, p. 86) กล่าวว่า การทดสอบด้วยสถิติ Chi – square นั้นถือว่าเป็น Measure of Badness – of – fit ดังนั้นในงานวิจัยครั้งนี้ผู้วิจัยจึงไม่ได้ศึกษาการทดสอบด้วยสถิติ Chi – square

1.2 ดัชนี GFI (Goodness of Fit Index)

GFI เป็นการวัดด้วยการเปรียบเทียบปริมาณของความแปรปรวนและความแปรปรวนร่วมของ S ที่ถูกทำนายโดย $\hat{\Sigma}$

$$GFI = 1 - \frac{F[S, \hat{\Sigma}]}{F[S, \Sigma(0)]}$$

ตัวเลขจะเป็นค่าน้อยที่สุดของ Fit Function เมื่อ โมเดลมีความสอดคล้อง ตัวส่วนจะเป็น Fit Function ของโมเดลอื่น ๆ ที่พบว่ามีความสอดคล้อง หรือเมื่อพารามิเตอร์ทั้งหมดเป็น 0 เขียนเป็นสูตรสำหรับการคำนวณได้ดังนี้

$$GFI = 1 - \frac{(S - \hat{\sigma})' W^{-1} (S - \hat{\sigma})}{S' W^{-1} S}$$

S คือ Variance – covariance Matrix ของกลุ่มตัวอย่าง

$\hat{\sigma}$ คือ Variance – covariance Matrix ของประชากรตามทฤษฎี

W คือ เมตริกซ์ถ่วงน้ำหนักที่ใช้ปรับค่าในการคำนวณ

ค่า GFI จะมีค่าระหว่าง 0 ถึง 1 ซึ่งมีเกณฑ์ที่ยอมรับได้แต่เดิม คือ มีค่าตั้งแต่ .90 ขึ้นไป (Kline, 2005, p. 145) แสดงว่าโมเดลมีความสอดคล้อง แต่ในปัจจุบันจะใช้ค่าตั้งแต่ .95 ขึ้นไป

1.3 ดัชนี AGFI (Adjusted Goodness of Fit Index)

AGFI คำนวณจากค่า GFI แต่พิจารณาถึงจำนวนตัวแปรที่วัดได้ และขนาดของกลุ่มตัวอย่างทั้งหมด AGFI จึงเป็นการปรับ df ของโมเดล ซึ่งแสดงปริมาณความแปรปรวนและความแปรปรวนร่วม ที่อธิบายได้ด้วยโมเดลโดยปรับแก้ด้วยของค่าความเป็นอิสระค่า AGFI จะมีค่าระหว่าง 0 ถึง 1 ซึ่งมีเกณฑ์ที่ยอมรับได้คือ มีค่าตั้งแต่ .90 ขึ้นไป (Kline, 2005, p. 145) มีสูตรดังนี้

$$AGFI = 1 - \frac{(p+q)(p+q+1)}{2df} (1 - GFI)$$

p คือ จำนวนตัวแปรสังเกตได้ (Observed Variables)

q คือ จำนวนตัวแปรทำนาย (Predictor Variables) การศึกษาความเป็นเอกมิติของโมเดลการวัดจะไม่มีกำหนดในโมเดล ดังนั้น $q = 0$

df คือ Degrees of Freedom ของโมเดล

1.4 การพิจารณาความคลาดเคลื่อน (Residual)

เป็นการพิจารณาความคลาดเคลื่อน หรือส่วนต่างระหว่างสมาชิกในเมตริกซ์ของ S และสมาชิกในเมตริกซ์ของ $\hat{\Sigma}$ ว่าเป็นเมตริกซ์ศูนย์ (Zero Matrix) หรือไม่ การเปรียบเทียบสมาชิกระหว่างเมตริกซ์ทั้งสองจะทำให้ได้เมตริกซ์ความคลาดเคลื่อน (Residual Matrix) เมื่อพบว่าส่วนต่างระหว่างสมาชิกมีค่าไม่เท่ากับ 0 หมายความว่า การกำหนดโมเดล (Model Specification) เกิดความคลาดเคลื่อน วิธีนี้จะเป็นวิธีง่ายที่สุด สำหรับการทดสอบความสอดคล้องของข้อมูลเชิงประจักษ์และโมเดลตามสมมติฐาน

เมตริกซ์ $\hat{\Sigma}$ ที่ใช้ในการทดสอบไม่สามารถนำมาจากประชากรทั้งหมด จึงต้องใช้เมตริกซ์จากกลุ่มตัวอย่าง S แทน การใช้ $S - \hat{\Sigma}$ จึงจึงมีลักษณะเดียวกับการทดสอบด้วยสถิติ Chi – square ค่า Residual เป็นบวก (+) หมายความว่า โมเดลที่สร้างขึ้นทำนายค่าความแปรปรวนได้ต่ำกว่าความแปรปรวนรวมของข้อมูลที่ได้จากกลุ่มตัวอย่าง ส่วนค่า Residual เป็นลบ (-) หมายความว่า โมเดลที่สร้างขึ้นทำนายค่าความแปรปรวนได้สูงกว่าความแปรปรวนรวมของข้อมูลที่ได้จากกลุ่มตัวอย่าง การใช้ความคลาดเคลื่อนเป็นค่าพิจารณาความสอดคล้องของข้อมูลและโมเดลที่สร้างขึ้น โดยความคลาดเคลื่อนใน SEM จะเป็นค่าความแปรปรวนรวมควรมีค่าใกล้ 0 และมีการแจกแจงแบบสมมาตร ถึงแม้ว่าผลการวิเคราะห์จะมีดัชนีหลายตัวที่แสดงค่าสอดคล้อง แต่มักพบว่าค่าความคลาดเคลื่อนประมาณ 1 – 2 ค่ามีค่าสูงมาก ในลักษณะนี้โมเดลจะมีความสอดคล้องเมื่อความคลาดเคลื่อนมีการแจกแจงแบบสมมาตรและมีค่าเฉลี่ยเข้าใกล้ศูนย์ หากการพิจารณาค่าความคลาดเคลื่อนนี้เองทำให้มีการพัฒนาดัชนีในการวัดความสอดคล้องของโมเดลต่อมาอีกหลายตัวด้วยกัน คือ

1.4.1 ดัชนี RMR (Root Mean Square Residual)

วิธีการใช้ความคลาดเคลื่อน จะใช้ค่าเฉลี่ยของความคลาดเคลื่อนโดยเป็นค่าเฉลี่ยของผลต่างของสมาชิกในครั้งของเส้นทแยงมุมและค่าผลต่างในแนวเส้นทแยงมุมของเมตริกซ์ ยกกำลังสองของผลต่างเพื่อไม่คิดเครื่องหมาย โมเดลที่มีความสอดคล้องควรมีค่าเฉลี่ยความคลาดเคลื่อนเข้าใกล้ศูนย์ RMR เป็นดัชนีที่วัดความคลาดเคลื่อนเฉลี่ยของข้อมูลจากกลุ่มตัวอย่างที่คลาดเคลื่อนไปจากโมเดลทางทฤษฎี (Average of the Fitted Residuals)

ความคลาดเคลื่อนของกลุ่มตัวอย่างที่ได้จากการคำนวณ มีผลมาจาก

ก. ความแตกต่างระหว่าง S และ $\hat{\Sigma}$

ข. มาตรฐานของตัวแปรสังเกตได้

ค. ความคลาดเคลื่อนจากการสุ่มตัวอย่าง

เมื่อตัวแปรที่มีมาตรวัดที่แตกต่างกัน ตัวแปรบางตัวใช้มาตรวัดที่มีพิสัยกว้าง (Large Range) จะบิดเบือนค่าเฉลี่ยของค่า Residual ทำให้ผลที่ได้ผิดไปด้วย จึงมีการทำให้ค่านี้เป็นคะแนนมาตรฐานเรียกว่า Root Standardized Mean Square Residual (SRMR) ถ้าค่า SRMR > |4.0| แสดงว่าโมเดลไม่สอดคล้องกับข้อมูลเชิงประจักษ์

ดัชนี RMR เป็นค่าเฉลี่ยของความคลาดเคลื่อนระหว่าง $S - \hat{\Sigma}$ ค่าที่น้อยแสดงถึงโมเดลสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ แต่ค่า RMR ขึ้นอยู่กับหน่วยการวัดของตัวแปร เมื่อตัวแปรมีสเกลการวัดที่ต่างกันมาก ตัวแปรบางตัวที่มีสเกลการวัดกว้างจะทำให้ค่าเฉลี่ยของ Residual บิดเบือนไป ทำให้ค่าที่ได้ผิดไปด้วย ดังนั้นอาจพิจารณาพร้อมกับค่าความคลาดเคลื่อนมาตรฐานของส่วนที่เหลือ (Standardized Residual) ซึ่งเป็นค่าของความคลาดเคลื่อนหารด้วยความคลาดเคลื่อนมาตรฐานของการประมาณค่า (Estimated Standard Error) ค่าคลาดเคลื่อนมาตรฐานไม่ควรเกิน |2.58| ซึ่งค่า Standardized RMR เป็นค่าสรุปของ Standardized Residual ควรมีค่าน้อยกว่า .05 จึงจะสรุปได้ว่า โมเดลสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์

1.4.2 ดัชนี RMSEA (Root Mean Square Error of Approximation) ค่ารากที่สองของค่าเฉลี่ยความคลาดเคลื่อนกำลังสองของการประมาณค่า ค่านี้ใช้ทดสอบสมมติฐาน $H_0: S = \hat{\Sigma}$ แต่ในทางความเป็นอิสระมาปรับแก้ โดยมีสูตรการคำนวณ ดังนี้ $RMSEA = (F_0/df)$ เมื่อ F_0 คือ Population Discrepancy Function Value หรือ ค่าฟังก์ชันความกลมกลืนเมื่อโมเดลสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ ถ้า F_0 เท่ากับศูนย์ RMSEA จะเท่ากับศูนย์ แสดงว่า โมเดลสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ดีมาก

ดัชนีนี้เป็นการวัดความแตกต่างต่อหน่วยขององศาความเป็นอิสระ (Discrepancy per Degrees of Freedom) โดย Brown and Cudeck (1993 cited in Jöreskog & Sorbom, 1993, p.124) เสนอให้อ่านค่า RMSEA ที่ .05 แสดงว่า มีความสอดคล้องมาก ถ้าค่าที่ได้สูงขึ้นถึง .08 แสดงว่าเกิดความคลาดเคลื่อนขึ้นในการประมาณค่าประชากร ส่วน Kline (2005, p. 139) เสนอว่าค่า RMSEA ที่ดีมาก ๆ ควรมีค่าน้อยกว่า .05 ค่าระหว่าง .05 – .08 หมายถึง โมเดลค่อนข้างสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ และค่ามากกว่า .10 แสดงว่า โมเดลยังไม่สอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ การอ่านค่า RMSEA ยังไม่สามารถระบุได้ชัดเจนว่า ถ้าค่าที่ได้แตกต่างจาก 0.05 เพียงเล็กน้อย จะยังคงถือว่ามีความสอดคล้องอยู่หรือไม่ จากการศึกษาจะพบว่า ค่า RMSEA เหมาะที่จะใช้ใน Confirmatory Model และ Competing Model เมื่อใช้กลุ่มตัวอย่างขนาดใหญ่ตั้งแต่ 500 ตัวอย่างขึ้นไป

1.4.3 ค่า ECVI (Expected Cross Validation Index) เป็นการทดสอบภาพรวมของความคลาดเคลื่อนระหว่าง S กับ $\hat{\Sigma}$ ถ้าโมเดลสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ ค่า ECVI ต้องมีค่าน้อยกว่าค่า ECVI for Saturated Model และน้อยกว่าค่า ECVI for Independence Model

1.4.4 ค่า NCP (Non – centrality Parameter) การทดสอบด้วยสถิติ Chi – square อาจปฏิเสธสมมติฐานศูนย์ได้ง่ายเนื่องจากข้อมูลมิได้แจกแจงแบบ Chi – square แต่มีการแจกแจงแบบ Non – central χ^2 (การแจกแจงแบบ Chi – square เป็นกรณีหนึ่งของการแจกแจงแบบ Non – central χ^2) ซึ่งมีค่า Non – centrality parameter เป็น λ โดยค่า λ จะแสดงความแตกต่างของเมตริกซ์ S กับ $\hat{\Sigma}$ ถ้า λ เท่ากับศูนย์แสดงว่าโมเดลสอดคล้องกับข้อมูลเชิงประจักษ์ ถ้าค่า λ ยิ่งมากยิ่งมีโอกาสปฏิเสธสมมติฐานศูนย์มาก โดยโปรแกรมจะแสดงค่า λ ในช่วงเชื่อมั่น 90% ถ้าโปรแกรมไม่แสดงหมายความว่าค่า λ ใหญ่มากจนไม่สามารถประมาณค่าช่วงความเชื่อมั่นได้

2. Incremental Fit Indices หรือ Model Comparison Fit Indices (ค่าดัชนี

ในการเปรียบเทียบโมเดล) (Bollen, 1989, pp. 269 – 276)

เมื่อทำการทดสอบด้วยสถิติ Chi – square ยังมีจุดอ่อนบางประการ นักวิชาการจึงหันมาสนใจ Fit Function (F) แทน ไม่ว่าจะเป็น F_{ML} หรือ F_{ULS} หรือ F_{GLS} ต่างก็เป็นฟังก์ชันของ S และ $\hat{\Sigma}$ ซึ่งใน F เองจะให้ค่า (Scalar) จำนวนหนึ่ง que แสดงถึงความแตกต่างระหว่าง S และ $\hat{\Sigma}$ มีคุณสมบัติบางประการดังนี้

- ก. ค่าของ F ต่ำสุดมีค่าเท่ากับ 0
 - ข. เมื่อสมมติฐานเป็นจริง ค่าการแจกแจงของ F จะมีความสัมพันธ์ในทางตรงกันข้ามกับค่า N และเมื่อ F เข้าใกล้ 0 ค่า N จะเข้าใกล้ ∞
 - ค. จากกลุ่มตัวอย่างและ โมเดลที่กำหนด หากเพิ่ม Free Parameter จะไม่มีผลต่อค่า F
- การเปรียบเทียบกับ โมเดลที่ตัวแปรไม่มีความสัมพันธ์กันเลย (Baseline Model) หรือ โมเดลอิสระ (Independence Model) ซึ่ง โมเดลที่นิยมใช้เป็น Baseline มากที่สุดคือ Null Model ที่กำหนดให้ตัวแปรสังเกตได้ทั้งหมดไม่มีความสัมพันธ์กัน เป็นการเปรียบเทียบค่า F ของ โมเดลตามทฤษฎีที่เคลื่อนตัวจากโมเดล Baseline จึงทำให้เรียกดัชนีที่ใช้บ่งชี้ความสอดคล้องของข้อมูลกับโมเดลในชุดนี้ว่า “Incremental Fit Index” ดัชนีในชุดนี้ประกอบด้วย

2.1 ดัชนี NFI (Normed Fit Index) ของ Bentler and Bonett (1980 cited in Bollen, 1989, p. 269)

$$NFI = \frac{F_b - F_m}{F_b} \quad \text{หรือ}$$

$$NFI = \frac{\chi_b^2 - \chi_m^2}{\chi_b^2}$$

F_b คือ Fit Function ของ Baseline Model

F_m คือ Fit Function ของข้อมูลกับโมเดลตามทฤษฎี

โดยที่ $(n-1)F_{ML}$ หรือ $(n-1)F_{ULS}$ เป็น χ^2 Estimator จึงนำค่า χ_b^2 แทน F_b และใช้ค่า χ_m^2 แทน F_m ค่า NFI มีค่าตั้งแต่ 0 ถึง 1 ถ้ายิ่งใกล้ 1 จะบอกถึงความสอดคล้องของข้อมูลกับโมเดลมากขึ้นเท่านั้น โดยข้อจำกัดของ NFI มีดังนี้

2.1.1 ไม่มีการควบคุม Degrees of Freedom (df) เหมือนกับค่า R^2 ใน Regression Analysis ที่มีค่า Adjust R^2 เป็นการปรับแก้ค่าที่มีผลมาจาก Degrees of Freedom ในโมเดลที่มีลักษณะซ้ำซ้อนมาก อาจมีค่า NFI สูง แม้ว่าจะมี df น้อยก็ตาม และมักจะเป็น Overfitting Data

2.1.2 ขนาดของกลุ่มตัวอย่าง (n) ไม่มีผลต่อค่าดัชนี แต่มีผลต่อ Sampling Distribution ของค่าดัชนี ซึ่งค่าดัชนีที่ไม่มีผลของ n ต่อ Sampling Distribution จะมีประโยชน์เมื่อใช้เปรียบเทียบโมเดลจากกลุ่มตัวอย่างที่มีขนาดไม่เท่ากัน

2.2 ดัชนี IFI (Incremental Fit Index) (Bollen, 1989, p. 271)

$$IFI = \frac{F_b - F_m}{F_b - \left(\frac{df_m}{n-1}\right)} \quad \text{หรือ}$$

$$IFI = \frac{\chi_b^2 - \chi_m^2}{\chi_b^2 - df_m}$$

การคำนวณ IFI จะมีผลจากขนาดกลุ่มตัวอย่าง เมื่อกลุ่มตัวอย่างขนาดเล็ก ค่า IFI จะมีค่ามากกว่ากลุ่มตัวอย่างขนาดใหญ่ การใช้ค่านี้จะใช้เมื่อเห็นว่าค่า NFI มีค่าน้อยลงเมื่อกลุ่มตัวอย่างขนาดเล็ก แต่ถ้าพิจารณาค่า IFI ของโมเดล 2 โมเดลที่มีค่า χ_b^2 และ χ_m^2 เดียวกัน โมเดลที่มีขนาดกลุ่มตัวอย่างที่เล็กกว่าจะมีค่า IFI ที่มากกว่า ซึ่ง ค่า IFI มีลักษณะดังนี้

2.2.1 จะมีค่าถึง 1 ด้วยขนาดกลุ่มตัวอย่างขนาดต่าง ๆ

2.2.2 ค่าที่ได้ไม่มีพิสัยที่แน่นอนว่าจะอยู่ระหว่าง 0 ถึง 1 หรือไม่

2.2.3 ค่าที่ได้มากกว่า 1 เมื่อเป็น Overfitting Data

เมื่อขนาดกลุ่มตัวอย่างใหญ่ขึ้น $[df_m / (n - 1)]$ จะเข้าใกล้ 0 ค่า IFI จะมีค่าใกล้เคียง

กับค่า NFI

2.3 ดัชนี RFI (Relative Fit Index) (Bollen, 1989, p. 272)

$$RFI = \frac{\frac{F_b}{df_b} - \frac{F_m}{df_m}}{\frac{F_b}{df_b}} \quad \text{หรือ}$$

$$RFI = \frac{\frac{\chi_b^2}{df_b} - \frac{\chi_m^2}{df_m}}{\frac{\chi_b^2}{df_b}}$$

ค่าที่ได้เกิดจากการหารค่า F_b และ F_m ด้วย df เป็นความสอดคล้องต่อหน่วยของ df ทั้งโมเดล Baseline และโมเดลที่ต้องการทดสอบ ค่า df อาจมีค่าคงที่หรือลดลงเมื่อโมเดลมีความซับซ้อนมากขึ้น ค่า RFI ในโมเดลที่มีจำนวน Parameter น้อย จะมีค่าน้อยกว่า ค่า RFI ของโมเดลที่มีความซับซ้อนมากขึ้น และมีค่าอยู่ระหว่าง 0 ถึง 1

RFI มีลักษณะเดียวกับ NFI ที่ไม่เปลี่ยนแปลงไปตามขนาดกลุ่มตัวอย่าง แต่ค่าเฉลี่ยของ Sampling Distribution จะเพิ่มขึ้นตามขนาดกลุ่มตัวอย่าง

2.4 ดัชนี TLI (Tucker – Lewis Index) หรือ ดัชนี NNFI (Non – normed Fit Index) ของ Tucker and Lewis และ Bonett and Bentler (1989 cited in Bollen, 1989, p. 273)

$$TLI/ NNFI = \frac{\frac{F_b}{df_b} - \frac{F_m}{df_m}}{\frac{F_b}{df_b} - \frac{1}{n-1}} \quad \text{หรือ}$$

$$TLI/ NNFI = \frac{\frac{\chi_b^2}{df_b} - \frac{\chi_m^2}{df_m}}{\frac{\chi_b^2}{df_b} - 1}$$

ดัชนีตัวนี้สร้างขึ้นเพื่อลดปัญหาเกี่ยวกับค่าเฉลี่ยของ Sampling Distribution ทั้ง RFI และ NNFI เป็นการแก้ df ของโมเดล Baseline ค่าของ TLI/ NNFI อาจมีค่าน้อยกว่า 0 หรือมีค่ามากกว่า 1 ได้ โมเดลจะมีความสอดคล้องเมื่อ TLI/ NNFI มีค่าตั้งแต่ .95 ขึ้นไป

2.5 ดัชนี CFI (Comparative Fit Index)

$$CFI = \frac{\chi_m^2 - df_m}{\chi_b^2 - df_b}$$

CFI เป็นดัชนีที่ปรับปรุงมาจาก NFI ของ Bentler & Bonett โดย CFI เป็น Normed ทำให้มีค่าอยู่ระหว่าง 0 ถึง 1 ความซับซ้อนของโมเดลไม่มีผลต่อดัชนีนี้ เมื่อดัชนีมีค่าตั้งแต่ .95 ขึ้นไป แสดงให้เห็นความสอดคล้องของโมเดล

ปัจจัยที่มีผลต่อดัชนีใน Incremental Fit Index

2.5.1 การเลือกโมเดล Baseline ที่มีข้อจำกัดมาก จะยิ่งทำให้โมเดลต้องมีการทดสอบความสอดคล้องมากขึ้น

2.5.2 การกำหนดมาตรฐานของค่าดัชนีไว้ล่วงหน้า เช่น จากงานวิจัยกำหนดค่าดัชนีไว้ .95 แต่ในรายงานใหม่ที่มีค่าดัชนี .85 หรือ .90 ทำให้เราไม่ยอมรับว่าโมเดลมีความสอดคล้อง แต่ถ้าในรายงานเดิมกำหนดค่าดัชนีไว้ .80 เมื่อรายงานใหม่มีค่าดัชนี .85 หรือ .90 จะยอมรับว่าโมเดลมีความสอดคล้อง

2.5.3 ความสัมพันธ์ระหว่างดัชนีความสอดคล้อง จากสูตรของดัชนีจะพบว่าค่าที่ได้จากแต่ละสูตรแตกต่างกัน ทำให้ต้องมีการกำหนดค่าจุดตัดที่แตกต่างกัน

2.5.4 การเลือกใช้วิธีการประมาณค่า F ที่แตกต่างกัน ทำให้ค่าของดัชนีที่ได้แตกต่างกันตามไปด้วย ไม่ว่าจะใช้วิธี ML หรือ ULS หรือ GLS

3. Parsimony Fit Indices (ค่าดัชนีวัดความประหยัดของโมเดล)

เป็นดัชนีที่อาศัยพื้นฐานมาจากการเปรียบเทียบโมเดลจำนวนหนึ่ง เป็นการพิจารณาโมเดลที่ดีกว่าโมเดลอื่น ๆ โดยอาศัยความซับซ้อนที่เพิ่มมากขึ้นในแต่ละโมเดล โมเดลที่ไม่ซับซ้อนจึงเป็นโมเดลที่มีเส้นทางที่ต้องประมาณค่าพารามิเตอร์จำนวนน้อย ดัชนีชุดนี้ ได้แก่ PNFI, CN, AIC และ PGFI

3.1 ดัชนี PNFI (Parsimony Normed Fit Index)

PNFI พัฒนามาจาก NFI โดยคูณด้วย PR โดยการพิจารณาค่า PNFI จะเช่นเดียวกับค่า NFI ซึ่งมีสูตรดังนี้

$$PNFI = \frac{df_m}{df_t} \quad \text{หรือ}$$

$$PNFI = \left(\frac{F_b - F_m}{F_b} \right) \left(\frac{df_m}{df_t} \right)$$

เมื่อ t คือ Total Degree of Freedom ของโมเดลทั้งหมด

3.2 Critical N (CN) ของ Heolter (1983 cited in Bollen, 1989, p. 277)

$$CN = \frac{\text{Critical } \chi^2}{F} + 1$$

χ^2 คือ ค่าวิกฤติของ χ^2 ที่มี df เท่ากับ โมเดลที่ต้องการทดสอบ และค่า α เท่ากับที่กำหนดไว้ และ F คือ ค่า F_{ML} หรือ F_{GLS} ของ S และ Σ

Bollen (1989, p. 277) เสนอให้ใช้จุดตัดของค่านี้ที่ $CN > 200$ ด้วยกลุ่มตัวอย่างขนาดใหญ่ การที่ N ไม่เกี่ยวข้องในสูตร ทำให้ค่า CN เท่ากันในทุกขนาดกลุ่มตัวอย่าง

3.3 ค่า Model AIC (Akaike's Information Criterion) เป็นการทดสอบภาพรวมของความคลาดเคลื่อนระหว่าง S กับ Σ ถ้าโมเดลสอดคล้องกลมกลืนกับข้อมูลเชิงประจักษ์ ค่า Model AIC ต้องน้อยกว่าค่า Saturated AIC และ Independence AIC นอกจากนี้ยังมีค่า Model CAIC (Consistent Version of AIC) ซึ่งเป็นค่า AIC ที่ปรับแก้ด้วยขนาดกลุ่มตัวอย่าง การแปลความหมายเหมือนค่า Model AIC

3.4 ดัชนี PGFI (Parsimony Goodness of Fit) คือ ดัชนีวัดความประหยัดของความเหมาะสมพอดี ที่แสดงปริมาณความแปรปรวนและความแปรปรวนร่วม ที่อธิบายได้ด้วยโมเดล โดยปรับแก้ด้วยความซับซ้อนของโมเดล ซึ่งมีเกณฑ์ที่ยอมรับได้คือ มีค่าไม่เกิน .50 ในดัชนีชุดนี้แม้ว่า AGFI จะไม่ขึ้นอยู่กับขนาดตัวอย่าง และจะมีค่าสูงสุดเมื่อ $S = \Sigma$ แต่การศึกษาของ Anderson and Gerbing (1984 cited in Bollen, 1989, p. 277) พบว่า AGFI จะมีค่าลดลงเมื่อเพิ่มจำนวนตัวแปรสังเกตได้ต่อตัวแปรคุณลักษณะแฝง โดยเฉพาะกลุ่มตัวอย่างขนาดเล็ก

ข้อพิจารณาเกี่ยวกับดัชนีวัดความสอดคล้องของโมเดล

1. ค่าที่ใช้ในการตัดสินความสอดคล้องของโมเดล เดิมนั้นนักวิจัยส่วนใหญ่จะยึดค่าตั้งแต่ .90 ขึ้นไป แต่ยังพบว่ายังคงเป็นโมเดลที่ไม่ถูกต้อง ทำให้มีการใช้ดัชนีในชุด Incremental Fit Index ตั้งแต่ .95 ขึ้นไป

2. ความซับซ้อนของโมเดลมีอิทธิพลต่อการวัดความสอดคล้องของโมเดล ไม่ว่าจะเป็นการมีตัวแปรสังเกตได้จำนวนมากต่อตัวแปรคุณลักษณะแฝง

3. การแจกแจงของข้อมูลมีอิทธิพลต่อการวัดความสอดคล้องของโมเดล และจะมีอิทธิพลต่อดัชนีใด Index มากกว่า Index Absolute Fit Index

การพิจารณาความสอดคล้องของโมเดลไม่ควรใช้ดัชนีชุดใดชุดหนึ่ง ดัชนีที่นิยมใช้ในการรายงานประกอบด้วยค่า χ^2 , df , CFI, NNFI และ RMSEA เมื่อโมเดลมีความซับซ้อนมากควรเสนอค่า PNFI ด้วย

Schemelleh, Moosbrugger and Müller (2003, pp. 23 – 27) ได้เสนอแนวทางในการตัดสินค่าดัชนีเป็น 2 ลักษณะ คือ ค่าที่แสดงความสอดคล้อง และค่าที่ยอมรับได้ว่ามีความสอดคล้อง สามารถแสดงได้ดังตารางที่ 2 – 6

ตารางที่ 2 – 6 เกณฑ์ดัชนีวัดความสอดคล้องกลมกลืนของโมเดล

Model Fit Indices	ค่าที่แสดงความสอดคล้อง ในระดับดี (Good Fit)	ค่าที่แสดงความสอดคล้อง ในระดับปานกลาง (Adequate Fit)
χ^2	$.05 < p \leq 1.00$	$.01 < p \leq .05$
χ^2/df	$0 < \chi^2/df \leq 2$	$2 < \chi^2/df \leq 3$
RMR	$0 \leq RMR \leq .05$	$.05 \leq RMR \leq .08$
SRMR	$0 \leq SRMR \leq .05$	$0 \leq SRMR \leq .05$
RMSEA	$0 \leq RMSEA \leq .05$	$.05 \leq RMSEA \leq .08$
NFI	$.95 \leq NFI \leq 1.00$	$.90 \leq NFI \leq .95$
TLI/ NNFI	$.97 \leq TLI \leq 1.00$	$.95 \leq TLI \leq .97$
CFI	$.97 \leq CFI \leq 1.00$	$.95 \leq CFI \leq .97$
GFI	$.95 \leq GFI \leq 1.00$	$.90 \leq GFI \leq .95$
AGFI	$.90 \leq AGFI \leq 1.00$	$.85 \leq AGFI \leq .90$

ผลกระทบจากการละเมิดข้อตกลงเบื้องต้นในการวิเคราะห์โมเดลสมการเชิงโครงสร้างและโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน

Jöreskog (1996 cited in Schumacker & Lomax, 2010, p. 19) ได้กำหนดเกณฑ์เกี่ยวกับตัวแปรที่มีระดับการวัดแบบเรียงอันดับ (Ordinal Scale) และแบบช่วง (Interval Scale) ในงานวิจัยว่า ตัวแปรที่มีสเกลการวัดไม่เกิน 15 Categories ในโปรแกรม PRELIS ถือว่ามีระดับการวัดแบบเรียงอันดับ (Ordinal Scale) และตัวแปรที่มีสเกลการวัดมากกว่า 15 Categories ถือว่าเป็นตัวแปรแบบต่อเนื่อง (Continuous Variable) ส่วน Muthén and Muthén (2002) กล่าวว่า ข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) เช่น มาตรวัดแบบ Likert ซึ่งพบมากในงานวิจัยทั่วไป โดยธรรมชาติของข้อมูลไม่สามารถนิยามการแจกแจงโค้งปกติ เพราะเป็นข้อมูลแบบไม่ต่อเนื่อง วิธีการประมาณค่าที่ใช้หลักการ Normal Theory ในการวิเคราะห์โมเดล SEM ส่วนมากมักจะให้ข้อสมมติ (Assumed) ว่าข้อมูลประชากรมีการแจกแจงแบบปกติพหุ (Multivariate Normal Distribution)

การตรวจสอบลักษณะการแจกแจงของข้อมูลก่อนดำเนินการในการวิเคราะห์โมเดลสมการเชิงโครงสร้างนั้นถือว่ามีความสำคัญยิ่ง ซึ่งจะช่วยในการเลือกใช้สถิติวิเคราะห์ที่มีความเหมาะสมกับลักษณะธรรมชาติของข้อมูล ให้ผลลัพธ์ที่มีความน่าเชื่อถือ ข้อตกลงเบื้องต้นสำหรับการใช้สถิติวิเคราะห์โดยเฉพาะ โมเดลสมการเชิงโครงสร้าง คือ การแจกแจงตัวแปรแบบปกติพหุ (Multivariate Normal) ซึ่งการตรวจสอบตัวแปรพหุนั้นยังไม่มีเกณฑ์ที่แน่ชัดในการกำหนด แต่ส่วนมากมักจะให้ข้อสมมติ (Assume) ว่ามีการแจกแจงแบบปกติพหุ (Multivariate Normality) ถ้า 1) การแจกแจงในแต่ละตัวแปร (Univariate) ทุกตัวแปรมีการแจกแจงแบบปกติ 2) การแจกแจงร่วมกันของตัวแปรแต่ละคู่มีการแจกแจงแบบปกติทวินาม (Bivariate Normality) และ 3) แผนภาพการกระจายของทุก ๆ ตัวแปรทวินาม (Bivariate) มีลักษณะเป็นเส้นตรง (Kline, 2005, pp. 48 – 49) สำหรับการตรวจสอบการแจกแจงในแต่ละตัวแปร (Univariate Distribution) ซึ่งมีลักษณะที่แตกต่างกัน 4 แบบ ได้แก่ การแจกแจงแบบเบ้ทางลบ (Negative Skewness) เบ้ทางบวก (Positive Skewness) การแจกแจงแบบโด่งสูงกว่าโค้งปกติ (Leptokurtic) โด่งต่ำกว่าโค้งปกติ (Platykurtic) (Kline, 2005, p. 49) ซึ่งการตรวจสอบสามารถคำนวณได้จากสูตรดังต่อไปนี้

$$\text{Skew Index} = \frac{S^3}{(S^2)^{3/2}} \quad \text{และ} \quad \text{Kurtosis Index} = \frac{S^4}{(S^2)^2}$$

$$\text{เมื่อ } S^2 = \frac{\sum (X - M)^2}{N}, S^3 = \frac{\sum (X - M)^3}{N}, S^4 = \frac{\sum (X - M)^4}{N}$$

ค่าลบหรือบวกของค่ามาตรฐานของความเบ้และความโค้ง จะแสดงถึงทิศทางและลักษณะการแจกแจง ค่า Skew Index = 0 หมายถึง มีการแจกแจงเป็นโค้งแบบสมมาตร และค่า Kurtosis Index = 3.00 หมายถึง มีการแจกแจงเป็นโค้งปกติ ถ้ามีค่ามากกว่า 3.00 หมายถึง มีการแจกแจงแบบโค้งสูงกว่าโค้งปกติ (Leptokurtic) ถ้ามีค่าน้อยกว่า 3.00 หมายถึง มีการแจกแจงแบบโค้งต่ำกว่าโค้งปกติ (Platyokurtic) และถ้าหากค่าสัมบูรณ์ของค่า Kurtosis Index มีค่ามากกว่า 8.00 ถือว่ามีค่าความโค้งแบบสุดโค้ง (Extreme Kurtosis) (Kline, 2005, p. 50)

Kline (2005, pp. 178 – 179) กล่าวว่า วิธีการประมาณค่าแบบ ML นั้นมีข้อดกลงเบื้องต้นเรื่องการแจกแจงแบบปกติพหุ (Multivariate Normal Distribution) ของตัวแปร แต่การวิเคราะห์ในโปรแกรมคอมพิวเตอร์ทั่วไปมักหลีกเลี่ยงข้อกำหนดนี้ เมื่อตัวแปรแบบต่อเนื่องมีการแจกแจงไม่เป็นโค้งปกติในระดับมาก (Severely Non – normal) การประมาณค่าแบบ ML นั้นจะทำให้ค่าความเอนเอียงของค่าประมาณพารามิเตอร์และค่าความคลาดเคลื่อนมาตรฐานเป็นลบสูง (Negative Biased) ค่าความคลาดเคลื่อนประเภทที่ 1 (Type I Error) และสถิติ Chi – square มีค่าสูงขึ้น แนวโน้มในการปฏิเสธโมเดลสมมติฐานมีอัตราเพิ่มขึ้นด้วย

Finney and DiStefano (2006) กล่าวว่า การละเมิดข้อดกลงเบื้องต้นเกี่ยวกับการแจกแจงปกติพหุ (Multivariate Normal Distribution) ของข้อมูลที่มีลักษณะจำแนกกลุ่ม (Categorical Data) จะให้ผลลัพธ์การประมาณค่าที่ไม่มีความเชื่อมั่น เช่น ค่าดัชนีวัดความสอดคล้องกลมกลืน มีค่าต่ำ มีความเอนเอียงทางลบของค่าประมาณพารามิเตอร์และค่าความคลาดเคลื่อนมาตรฐาน ระดับความน่าจะเป็นที่ผลลัพธ์จากข้อมูลเชิงประจักษ์จะไม่สอดคล้องกับ โมเดลสมมติฐานนั้นขึ้นอยู่กับระดับการแจกแจงที่ไม่เป็นโค้งปกติ (Degree of Non – normality) และจำนวนกลุ่มที่จำแนก (Number of Categories) (DiStefano, 2002; Muthén & Kaplan, 1985; Myers, Ahn, & Jin 2011) ในกรณีที่จำนวนกลุ่มที่จำแนกน้อยกว่า 5 Categories นั้น Finney and DiStefano (2006) แนะนำให้ใช้วิธีประมาณค่าพารามิเตอร์แบบ Categorical Variable Methodology: CVM หรือวิธีการ WLSMV (Weighted Least Squares with Mean and Variance Adjusted Estimation)

ในเทคนิคทางสถิติส่วนใหญ่ โดยเฉพาะสมการโครงสร้างเชิงเส้น ถ้าวิธีการประมาณค่าพารามิเตอร์นั้นจะต้องเป็นไปตามข้อดกลงเบื้องต้นจึงจะทำให้ผลการวิจัยนั้นน่าเชื่อถือ ซึ่งผลลัพธ์ที่ต้องการศึกษา ได้แก่ ค่าพารามิเตอร์ (Parameter Values) ค่าความคลาดเคลื่อนมาตรฐาน (Standard Error) และ ดัชนีวัดความสอดคล้องกลมกลืน (Fit Indices) ซึ่งวิธีการประมาณค่าที่อาศัยหลักของทฤษฎีการแจกแจงปกติ (Normal Theory) ที่ใช้กันโดยทั่วไปสองวิธี คือ ML

(Maximum Likelihood) และ วิธี GLS (General Least Squares) ซึ่งหากข้อมูลเป็นไปตามข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงปกติพหุแล้ว ตัวแปรจะมีลักษณะเป็นข้อมูลแบบต่อเนื่องโดยธรรมชาติอยู่แล้ว แต่สำหรับข้อมูลที่มีลักษณะจัดกลุ่ม (Categorical Data) เช่น ตัวแปรทวิภาค (Dichotomies) หรือแม้แต่มาตราวัดแบบ Likert นั้น ไม่สามารถนิยามเกี่ยวกับทฤษฎีการแจกแจงปกติได้ เพราะโดยธรรมชาติของข้อมูลมีลักษณะเป็นแบบไม่ต่อเนื่อง (Discrete Data) (Kaplan, 2000 cited in Finney & DiStefano, 2006, p. 271) ดังนั้น หลายครั้งที่ตัวประมาณค่าในทฤษฎีการแจกแจงปกติ (NT Estimators) มีข้อตกลงเกี่ยวกับลักษณะของตัวแปรสังเกตได้จะต้องเป็นข้อมูลแบบต่อเนื่องที่มีการแจกแจงปกติ ซึ่งถ้าหากตัวประมาณค่าในทฤษฎีการแจกแจงปกติเป็นไปตามข้อตกลงเบื้องต้นจะมีผลทำให้การประมาณค่าพารามิเตอร์ประกอบด้วยคุณสมบัติ ดังนี้ (Finney & DiStefano, 2006, pp. 270 – 271)

1. ปราศจากความเอนเอียง (Unbiasedness)
2. ความมีประสิทธิภาพ (Efficiency)
3. ความคงเส้นคงวา (Consistency)

โดยทั่วไปการประมาณค่าพารามิเตอร์ด้วยวิธี ML จะได้รับผลกระทบจากการแจกแจงของข้อมูลที่ไม่เป็นโค้งปกติขึ้นอยู่กับระดับของความไม่เป็นโค้งปกติของการแจกแจง นักวิจัยต้องพิจารณาตรวจสอบการแจกแจงข้อมูลของตัวแปรสังเกตได้ก่อนดำเนินการวิเคราะห์ข้อมูลและเลือกวิธีประมาณค่าพารามิเตอร์ ซึ่งมีดังนี้ 3 ข้อที่บ่งบอกถึงการแจกแจงของข้อมูลที่ไม่เป็นโค้งปกติ ได้แก่

1. ค่าความเบ้ของการแจกแจงแต่ละตัวแปร (Univariate Skew)
2. ค่าความโด่งของการแจกแจงแต่ละตัวแปร (Univariate Kurtosis)
3. ค่าความโด่งของการแจกแจงพหุตัวแปร (Multivariate Kurtosis)

ในการวิจัยทางสังคมศาสตร์ ข้อมูลที่มีลักษณะการแจกแจงของข้อมูลที่ไม่เป็นโค้งปกติและเป็นข้อมูลลักษณะจำแนกกลุ่ม (Categorical Data) มักพบอยู่บ่อยครั้ง ดังนั้นจึงมีคำถามเกี่ยวกับความแกร่งของวิธีการประมาณค่าแบบ ML ภายใต้เงื่อนไขดังกล่าว จากงานวิจัยที่ผ่านมาได้ศึกษาผลกระทบจากการแจกแจงของข้อมูลที่ไม่เป็นโค้งปกติที่มีต่อค่าต่าง ๆ ได้แก่ ค่าสถิติ Chi – square (χ^2) ค่าดัชนีวัดระดับความกลมกลืน (Fit Indices) ค่าประมาณพารามิเตอร์ (Parameter Estimates) และ ค่าคลาดเคลื่อนมาตรฐาน (Standard Error) ซึ่งพบว่า วิธีการประมาณค่าแบบ ML ค่าสถิติต่าง ๆ จะเกิดความเอนเอียงมากขึ้นเมื่อระดับการแจกแจงของข้อมูลที่ไม่เป็นโค้งปกติเพิ่มมากขึ้น (Bollen, 1989; Chou, Bentler, & Storra, 1991; Finney & DiStefano, 2006, p. 273)

ถ้ามีการกำหนดลักษณะเฉพาะของโมเดลอย่างถูกต้องแล้ว แต่มีการละเมิดข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงแบบปกติพหุ จะพบว่า วิธีการ ML จะให้ค่า Chi – square ที่สูงขึ้น เมื่อระดับการแจกแจงของข้อมูลที่ไม่เป็นโค้งปกติเพิ่มมากขึ้น (Chou et al., 1991; Curran et al., 1996; Hu, Bentler & Kano, 1992 cited in Finney & DiStefano, 2006, p. 273) ค่าความโค้ง โดยเฉพาะการแจกแจงที่มีความโค้งสูงกว่าโค้งปกติ (Leptokurtic) หรือค่าความโค้งเป็นบวกสูง จะส่งผลต่อค่า Chi – square มากที่สุด การที่ค่า Chi – square มีค่าสูงมากจะส่งผลให้เกิดความคลาดเคลื่อนประเภทที่ 1 (Type I Error) เพิ่มมากขึ้น และทำให้มีโอกาสปฏิเสธโมเดลสมมติฐานสูงขึ้น นอกเหนือจากค่า Chi – square แล้ว ยังส่งผลต่อค่าดัชนีวัดระดับความกลมกลืนต่าง ๆ ด้วย

Hu and Bentler (1998, p. 427 cited in Finney & DiStefano, 2006, p. 273) กล่าวว่า ค่าดัชนีวัดระดับความกลมกลืนจะแสดงผลลัพธ์ที่ดีกว่า ถ้าการทดสอบ Chi – square นั้นแสดงผลลัพธ์ที่ดีจากผลการวิจัยที่ผ่านมา พบว่า ถ้าข้อมูลมีการแจกแจงที่ไม่เป็นโค้งปกติในระดับปานกลางถึงมาก และกลุ่มตัวอย่างมีขนาดเล็ก ($n \leq 250$) วิธีการ ML จะให้ค่าดัชนีวัดระดับความกลมกลืน เช่น TLI, CFI หรือ RMSEA มีแนวโน้มที่จะปฏิเสธโมเดลสมมติฐานสูงขึ้น

Bollen (1989, p. 433 cited in Finney & DiStefano, 2006, p. 274) กล่าวว่า งานวิจัยส่วนใหญ่ที่ใช้การประมาณค่าพารามิเตอร์แบบ ML นั้นมักใช้ในเงื่อนไขที่ข้อมูลตัวแปรสังเกตได้เป็นข้อมูลแบบต่อเนื่องและมีการแจกแจงปกติพหุ แต่ในงานวิจัยทางสังคมศาสตร์มักพบว่าข้อมูลไม่ได้มีลักษณะดังกล่าว บ่อยครั้งที่นักวิจัยมีการเก็บรวบรวมข้อมูลและดำเนินการกับข้อมูลที่มีระดับการวัดแบบเรียงอันดับ (Ordered Scale) โดยเฉพาะมาตรวัดของ Likert ซึ่งนักวิจัยมักจะปรับ (Treat) ให้เป็นข้อมูลแบบต่อเนื่อง ซึ่งมาตรวัดเรียงอันดับ (Ordinal Scale) นั้นถือว่าเป็นระดับการวัดที่มีประสิทธิภาพต่ำถึงแม้ว่าจะมีการแจกแจงแบบปกติ เพราะโดยธรรมชาติของข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) มีลักษณะเป็นข้อมูลแบบไม่ต่อเนื่องและไม่สามารถนิยามการแจกแจงแบบปกติ

Bollen (1989 cited in Finney & DiStefano, 2006, p. 276) กล่าวถึงการสร้างโมเดลการวิจัยจากข้อมูลที่มีมาตรวัดเรียงอันดับว่า บ่อยครั้งที่นักวิจัยมักจะละเลยการพิจารณาความเป็นธรรมชาติที่เป็นข้อมูลแบบจำแนกกลุ่ม (Categorical Data) และประยุกต์ใช้การประมาณค่าพารามิเตอร์ด้วยวิธี ML โดยการสร้างเมตริกซ์ความแปรปรวน – ความแปรปรวนร่วม ภายใต้เงื่อนไขการคำนวณด้วยเมตริกซ์สหสัมพันธ์แบบเพียร์สัน (PPM) ในเทคนิคนี้นักวิจัยจะปรับ (Treat) ข้อมูลแบบเรียงอันดับเป็นข้อมูลแบบต่อเนื่อง โดยเชื่อว่าข้อมูลแบบเรียงอันดับที่ความถี่ช่วงของอันดับที่ละเอียดขึ้นน่าจะช่วยให้เข้าใจลักษณะเป็นข้อมูลแบบต่อเนื่องโดยประมาณ และค่าสหสัมพันธ์ก็จะเข้าใจค่าพารามิเตอร์จริงด้วย ถ้าจำนวนกลุ่มในการจำแนกกลุ่มน้อย ๆ (Fewer

Categories) จะมีผลทำให้ค่าสหสัมพันธ์ PPM มีประสิทธิภาพต่ำและการประมาณค่าไม่เข้าใกล้ค่าจริงมากขึ้น ซึ่งส่งผลกระทบต่อค่าสถิติ Chi – square และดัชนีวัดระดับความกลมกลืน เพราะโดยทั่วไปดัชนีวัดระดับความกลมกลืนจะให้ผลลัพธ์ที่ดีถ้าการประมาณค่าอยู่ภายใต้เงื่อนไขการแจกแจงแบบปกติโดยประมาณและเมื่อข้อมูลแบบเรียงอันดับ 5 ระดับ ถูกปรับเป็นข้อมูลแบบต่อเนื่อง ค่า Chi – square มีประสิทธิภาพค่อนข้างดีเมื่อข้อมูลเป็นแบบจำแนกกลุ่ม 4 กลุ่ม แต่ถ้าน้อยกว่า 4 กลุ่ม พบว่า จะทำให้สถิติ Chi – square มีค่าสูงขึ้น ส่วนดัชนีวัดความกลมกลืน เช่น GFI, AGFI, RMR จะมีค่าน้อยลงเมื่อกลุ่มตัวอย่างมีขนาดเล็ก และเมื่อข้อมูลแบบเรียงอันดับ ถูกปรับเป็นข้อมูลแบบต่อเนื่องจะไม่ส่งผลกระทบมากนัก เมื่อเป็นข้อมูลแบบเรียงอันดับ อย่างน้อย 5 ระดับ (Bollen, 1989; Muthén & Kaplan, 1985; Babakus, Ferguson, & Jöreskog, 1987; Hutchinson & Olmos, 1998; Green et al., 1997; Finney & DiStefano, 2006, p. 276) นอกจากนี้ยังส่งผลกระทบต่อค่าพารามิเตอร์จากการประมาณและค่าส่วนเบี่ยงเบนมาตรฐาน เช่น ในงานวิจัยที่ผ่านมา ภายใต้เงื่อนไขต่อไปนี้ 1) ข้อมูลเป็นแบบเรียงอันดับอย่างน้อย 5 ระดับ 2) ข้อมูลมีการแจกแจงแบบปกติโดยประมาณ 3) ข้อมูลถูกปรับเป็นข้อมูลแบบต่อเนื่อง และ 4) ใช้การประมาณค่าพารามิเตอร์ด้วยวิธี ML ผลการศึกษา พบว่า การประมาณค่าพารามิเตอร์และค่าความสัมพันธ์ขององค์ประกอบ (Factor Correlation) มีประสิทธิภาพค่อนข้างต่ำ ค่าเบี่ยงเบนมาตรฐานจะมีความอ่อนไหวต่อความเป็นข้อมูลแบบจำแนกกลุ่มมากกว่า ซึ่งพบว่า ทำให้เกิดความเอนเอียงแบบลบ (Negative Bias) (Babakus et al., 1987; Muthén & Kaplan, 1985; West, Finch, & Curran, 1995; Finney & DiStefano, 2006, p. 277) ถ้าค่าเบี่ยงเบนมาตรฐานมีค่าต่ำ จะทำให้การทดสอบนัยสำคัญทางสถิติของค่าประมาณพารามิเตอร์มีโอกาสเกิดความคลาดเคลื่อนประเภท 1 (Type I Error) มากขึ้น เมื่อจำนวนกลุ่มจำแนก (Categories) ลดลง จะทำให้การประมาณค่าพารามิเตอร์และค่าเบี่ยงเบนมาตรฐานมีประสิทธิภาพต่ำลงมาก แม้ว่าจะมีการแจกแจงข้อมูลแบบสมมาตร

Muthén and Kaplan (1985) กล่าวว่า โดยปกติแล้วข้อมูลที่มีมาตรวัดเรียงอันดับ จะไม่นิยามเกี่ยวกับการแจกแจงแบบปกติ อย่างไรก็ตามถ้าตัวแปรสังเกตได้ที่เป็นตัวแปรแบบจัดกลุ่มมีหลาย ๆ กลุ่ม (อย่างน้อย 5 ระดับขึ้นไป) และมีการแจกแจงแบบปกติ การประมาณค่าพารามิเตอร์ด้วยวิธี ML นั้นจะไม่ส่งผลกระทบมากต่อความเอนเอียง (Bias) ในดัชนีวัดระดับความกลมกลืน ค่าประมาณพารามิเตอร์ และค่าเบี่ยงเบนมาตรฐาน แต่เทคนิคการประมาณค่าแบบนี้ จะมีความแม่นยำน้อยลง เมื่อข้อมูลตัวแปรสังเกตได้มีการแจกแจงไม่เป็นโค้งปกติ และจำนวนการจำแนกกลุ่มลดลง (Number of Categories Decrease) ซึ่งมีผลทำให้ค่าสหสัมพันธ์แบบเพียร์สัน (PPM) มีประสิทธิภาพน้อยลงเช่นกัน

Finney and DiStefano (2006, p. 277) กล่าวว่า การประมาณค่าพารามิเตอร์ด้วยวิธี ML ภายใต้งैนไข การแจกแจงข้อมูลไม่เป็น โค้งปกติและลักษณะข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical) จะส่งผลต่อสถิติ Chi – square และดัชนี RMR (Root Mean Squared Residual) ให้มีค่าสูงขึ้น และค่า GFI, NNFI, CFI จะมีประสิทธิภาพในการประมาณค่าต่ำลง หมายความว่า โมเดลสมมติฐานไม่สอดคล้องกับข้อมูลเชิงประจักษ์ จึงควรพิจารณาถึงเหตุผลที่จะหลีกเลี่ยงการใช้เทคนิคการประมาณค่านี้ นอกจากนี้ยังส่งผลกระทบต่อค่าพารามิเตอร์จากการประมาณและค่าเบี่ยงเบนมาตรฐาน เพราะเมื่อตัวแปรสังเกตได้เป็นข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical) มีการแจกแจงในแต่ละตัวแปรแบบเบ้ (Univariate Skewness) และการแจกแจงในแต่ละตัวแปรแบบ โค้ง (Univariate Kurtosis) ในระดับที่มากขึ้น การประมาณค่าพารามิเตอร์ด้วยวิธี ML จะมีผลทำให้เกิดความเอนเอียงทางลบ (Negative Bias) ในค่าประมาณพารามิเตอร์และค่าเบี่ยงเบนมาตรฐานมากขึ้นด้วย ระดับความเอนเอียงจะเพิ่มมากขึ้นเมื่อกลุ่มตัวอย่างมีขนาดเล็กลง จำนวนกลุ่ม (Categories) ลดลง ค่าความสัมพันธ์ระหว่างองค์ประกอบหรือตัวแปรสังเกตได้ต่ำลง เมื่อระดับการแจกแจงข้อมูลไม่เป็น โค้งปกติมากขึ้น (Babakus et al., 1987; Muthén & Kaplan, 1985)

งานวิจัยที่เกี่ยวข้อง

Muthén and Kaplan (1985) ได้ศึกษาเปรียบเทียบการประมาณค่าพารามิเตอร์ 4 วิธี ได้แก่ ML (Maximum Likelihood), GLS (Generalized Least Square), ADF (Asymptotically Distribution Free) และ CMV (Categorical Method Variables) ในโมเดล CFA ด้วยการจำลองข้อมูลประชากรที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Variables) 5 ระดับ โดยกำหนดโมเดล CFA องค์ประกอบเดียวที่มีตัวแปรสังเกตได้ 4 ตัวแปร กลุ่มตัวอย่างถูกจำลองจากข้อมูลประชากรและจำแนกตามลักษณะการแจกแจงที่ต่างกัน 5 แบบ (มีการแจกแจงเป็น โค้งปกติ ไปจนถึงมีการแจกแจงแบบเบ้ระดับมาก รวมถึงลักษณะการแจกแจงแบบ โค้งมากในกรณีสมมาตร) ซึ่งผลการศึกษา พบว่า ในกรณีตัวแปรมีการแจกแจงเบ้ระดับมาก ค่าสถิติ Chi – square จากการทดสอบด้วยวิธี ML และ GLS จะมีค่าสูง ในขณะที่วิธี ADF นั้นค่า Chi – square มีความแรงมากกว่า ในกรณีตัวแปรมีการแจกแจงแบบ โค้งระดับมาก การประมาณค่าด้วยวิธี ML และ GLS สำหรับข้อมูลลักษณะจำแนกกลุ่มแบบเรียงอันดับนั้น ไม่ส่งผลกระทบต่อดัชนีความสอดคล้องของโมเดล ส่วนข้อจำกัดของการศึกษานี้คือ ไม่ได้ศึกษาผลกระทบต่อค่าประมาณพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐาน ประชากรมีขนาดเล็ก ($N = 1,000$) และการศึกษาในแต่ละเงื่อนไขทำการศึกษาชุดข้อมูลเพียง 25 การทดลองซ้ำ (Replications)

Babakus, Ferguson and Jöreskog (1987) ได้ศึกษาความอ่อนไหว (Sensitive) ของวิธีการประมาณค่าแบบ Maximum Likelihood (ML) ในเงื่อนไขละเมิดข้อตกลงเบื้องต้นเกี่ยวกับมาตรวัด (Measurement Scale) และลักษณะการแจกแจงของข้อมูล โดยการเปรียบเทียบการใช้เมตริกซ์สหสัมพันธ์ 4 แบบ คือ Product – moment, Polychoric, Spearman's Roh และ Kendall's Tau – b ในโมเดล CFA องค์ประกอบเดียว ตัวแปรสังเกตได้ 4 ตัวแปร ทำการจำลองข้อมูลตัวแปรจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) 5 ระดับ ในแต่ละเงื่อนไขใช้ข้อมูลจากการจำลอง 300 การทดลองซ้ำ (Replications) และทำการศึกษาประสิทธิภาพการประมาณค่าพารามิเตอร์ ได้แก่ Convergence Rate, Improper Solution, Accuracy of the Loading Estimates, Fit Statistics และ Estimated Standard Errors ผลการศึกษา พบว่า การใช้เมตริกซ์สหสัมพันธ์โพลีคลอริก (Polychoric) จะให้ค่าประมาณน้ำหนักองค์ประกอบของที่มีความเหมาะสมมากกว่าเมตริกซ์สหสัมพันธ์แบบอื่น ๆ ส่วนดัชนีวัดความสอดคล้องกลมกลืนให้ค่าไม่เหมาะสมในทุก ๆ เงื่อนไข ค่าความคลาดเคลื่อนมาตรฐานจากการประมาณค่าให้ผล Overestimates ในทุกเงื่อนไข และเมื่อข้อมูลมีการแจกแจงแบบเบ้มากขึ้นจะทำให้การประมาณค่ามีประสิทธิภาพแยกลงสรุปโดยภาพรวม พบว่า การประมาณค่าด้วยวิธีการ ML ในเงื่อนไขข้อมูลที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) และการแจกแจงของข้อมูลไม่เป็นโค้งปกติ (Non – normality) นั้น การใช้เมตริกซ์สหสัมพันธ์โพลีคลอริกมีความเหมาะสมมากกว่าเมตริกซ์สหสัมพันธ์แบบเพียร์สัน (PPM) และเมตริกซ์สหสัมพันธ์แบบอื่น ๆ

Rigdon and Ferguson (1991) ได้ศึกษาผลลัพธ์การประมาณค่าโมเดล ได้แก่ Nonconvergence, Improper Solution, ค่าประมาณพารามิเตอร์, ค่าคลาดเคลื่อนมาตรฐาน และดัชนีวัดความกลมกลืนต่าง ๆ ภายใต้การจำลองข้อมูลประชากรโมเดล CFA 2 องค์ประกอบ องค์ประกอบละ 4 ตัวแปรสังเกตได้ มีลักษณะเป็นข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) 5 ระดับ ซึ่งมีการแจกแจงเป็นโค้งปกติโดยประมาณ และกลุ่มตัวอย่างได้จากการจำลองจากโมเดลประชากร ซึ่งมีการแจกแจงในระดับต่าง ๆ (โค้งปกติ ถึง เบ้ระดับมาก) ขนาดกลุ่มตัวอย่างต่างกัน 3 ขนาด และวิธีการประมาณค่าที่ต่างกัน 4 วิธี คือ ML (Maximum Likelihood), GLS (Generalized Least Square), WLS (Weighted Least Square) และ DWLS (Diagonally Weighted Least Square) โดยในแต่ละเงื่อนไขทำการจำลองชุดข้อมูลเงื่อนไขละ 300 การทดลองซ้ำ ผลการศึกษา พบว่า ค่าประมาณของพารามิเตอร์ในทุกวิธีการประมาณค่าและทุกขนาดกลุ่มตัวอย่างแสดงค่าความเอนเอียงเล็กน้อย ส่วนค่าคลาดเคลื่อนมาตรฐานแสดงค่าความเอนเอียงในวิธี WLS และพบว่าค่าสถิติ Chi – square มีค่าสูงมากในวิธีการ ML, GLS และ DWLS ส่วนวิธี WLS มีค่าสถิติ Chi – square ที่ต่ำกว่าค่า Expected Values ส่วนอัตรา

Nonconvergency และ Improper Solution มีความแรงต่อผลกระทบจากขนาดกลุ่มตัวอย่าง ลักษณะการแจกแจงและวิธีการประมาณค่า แต่เมื่อการแจกแจงแบบเบ้ระดับมาก จะพบอัตรา Nonconvergency และ Improper Solution มากขึ้นด้วย ข้อจำกัดของการศึกษานี้ คือ การจำลองข้อมูล บนฐานของโมเดล CFA เพียงโมเดลเดียว และข้อมูลมีลักษณะจำแนกกลุ่มแบบเรียงอันดับ ซึ่งในงานวิจัยที่ผ่านมา พบว่า ประเภทของโมเดลมีผลต่อการแสดงผลลัพธ์ต่าง ๆ ซึ่งนักวิจัย ควรพิจารณาเปรียบเทียบโครงสร้างของโมเดลที่แตกต่างกันด้วย

Chou, Bentler and Satorra (1991) ได้ศึกษาการแสดงผลของสถิติทดสอบ Chi – square ด้วยวิธีการประมาณค่าที่แตกต่างกัน 4 วิธี คือ SB – χ^2 , ML – χ^2 , Robust ML – χ^2 และ ADF – χ^2 ในโมเดล CFA สององค์ประกอบ องค์ประกอบละ 3 ตัวแปรสังเกตได้ ภายใต้เงื่อนไขการแจกแจงไม่เป็นโค้งปกติในระดับต่าง ๆ 6 ระดับ (ค่าความเบ้ ตั้งแต่ -2 ถึง 2 และค่าความโด่ง ตั้งแต่ 1 ถึง 8) และขนาดกลุ่มตัวอย่างที่ต่างกัน (200 และ 400) ในแต่ละเงื่อนไขศึกษาชุดข้อมูล เงื่อนไขละ 100 การทดลองซ้ำ (Replications) ผลการศึกษา พบว่า SB – χ^2 แสดงผลได้ดีกว่า ML – χ^2 และ ADF – χ^2 ให้ค่าที่แย่ที่สุด ค่าเบี่ยงเบนมาตรฐานจากวิธี Robust ML และ ADF แสดงผลได้ดีกว่าวิธีการ ML แต่วิธี ML มีความแรงและแสดงผลได้ดีกว่าวิธีการอื่น ๆ ภายใต้เงื่อนไขการแจกแจงแบบสมมาตรและโค้งแบนราบ (Platykurtic) และแบบไม่สมมาตรที่ความโด่ง เป็นศูนย์ และค่าเบี่ยงเบนมาตรฐานแสดงค่าความเอนเอียงในระดับเล็กน้อยในทุก ๆ วิธีการ ถึงแม้ว่าค่าเบี่ยงเบนมาตรฐานจากวิธี ML จะค่อนข้างมีความแรงต่อการฝ่าฝืนข้อตกลงเรื่องการแจกแจง และกลุ่มตัวอย่างขนาดเล็ก แต่พบว่า เมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้นค่าเบี่ยงเบนมาตรฐาน จากวิธี Robust ML และ ADF จะแสดงผลได้เหมาะสมมากกว่า

Muthén and Kaplan (1992) ได้ปรับปรุงและแก้ไขการศึกษาเดิม Muthén and Kaplan (1985) โดยเพิ่มการศึกษาจำนวนการทดลองซ้ำ (Replications) ขนาดกลุ่มตัวอย่าง และ โมเดล ประชากรที่แตกต่างกัน โดยตัวแปรอิสระที่ศึกษา คือ 1) ขนาดโมเดล CFA ที่ต่างกัน 4 ลักษณะ ได้แก่ โมเดล 2 องค์ประกอบ 6 ตัวแปรสังเกตได้, โมเดล 3 องค์ประกอบ 9 ตัวแปรสังเกตได้, โมเดล 3 องค์ประกอบ 12 ตัวแปรสังเกตได้ และ โมเดล 3 องค์ประกอบ 15 ตัวแปรสังเกตได้ 2) ลักษณะการแจกแจงข้อมูลตัวแปรที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ 5 แบบ (ตั้งแต่โค้งปกติจนถึงเบ้ระดับมาก) และ 3) วิธีการประมาณค่า 2 วิธี คือ GLS และ ADF ผลการศึกษา พบว่า เมื่อขนาดของโมเดลใหญ่ขึ้นและค่าความเบ้ของการแจกแจงมากขึ้น ค่าสถิติ Chi – square จะมีค่าสูงขึ้น ค่า ADF – χ^2 ค่อนข้างมีความแรงต่อการแจกแจงที่ไม่เป็นโค้งปกติและโมเดลขนาดเล็ก แต่อย่างไรก็ตามเมื่อขนาดของโมเดลใหญ่ขึ้นค่า ADF – χ^2 ก็จะมีค่าสูงขึ้น ในกรณีที่กลุ่มตัวอย่างขนาดเล็ก ($n < 500$) ค่า GLS – χ^2 และ ADF – χ^2 จะมีค่าสูงขึ้น ในการแจกแจงที่มีลักษณะเบ้

มากขึ้น พบว่า ค่าคลาดเคลื่อนมาตรฐานมีอำนาจการประมาณค่าต่ำ (Underestimated) และมีความเอนเอียงเพิ่มมากขึ้น ส่วนค่าคลาดเคลื่อนมาตรฐานจากวิธี ADF จะมีความเอนเอียงเพิ่มมากขึ้น เมื่อโมเดลมีขนาดใหญ่ขึ้นและกลุ่มตัวอย่างเล็กลง

Hu, Bentler and Kano (1992) ได้ใช้การจำลองข้อมูลในการศึกษาเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์แบบต่าง ๆ คือ ML, GLS, ADF และ SB - χ^2 ในโมเดล CFA 3 องค์ประกอบ ตัวแปรสังเกตได้ 15 ตัวแปร โดยมีเงื่อนไขการละเมิดข้อตกลงเบื้องต้นใน 3 ลักษณะ คือ ข้อมูลมีการแจกแจงเป็นโค้งปกติ, ความเป็นอิสระของตัวแปร และกลุ่มตัวอย่างขนาดใหญ่ การสุ่มตัวอย่างจากโมเดลประชากรด้วยขนาดกลุ่มตัวอย่างที่แตกต่างกัน ผลการศึกษาพบว่า ค่าสถิติทดสอบ Chi - square ในเงื่อนไขการแจกแจงของตัวแปรแฝงและตัวแปรสังเกตได้ ที่อิสระต่อกัน ซึ่งมีการแจกแจงที่ไม่เป็น โค้งปกติ (Non - normal) และกลุ่มตัวอย่างขนาดใหญ่ ค่า ML - χ^2 , GLS - χ^2 จะให้ผลลัพธ์ที่ต่ำกว่า ADF - χ^2 ส่วนในเงื่อนไขการแจกแจงตัวแปรแฝงที่ไม่เป็นอิสระต่อกัน ADF - χ^2 จะให้ผลลัพธ์ที่ต่ำกว่า ML - χ^2 , GLS - χ^2 ภายใต้เงื่อนไขกลุ่มตัวอย่างขนาดกลาง วิธีการ ADF มีแนวโน้มจะให้ค่าประมาณที่ดีกว่าวิธีการอื่น ๆ แต่เมื่อตัวแปรมีการแจกแจงแบบโค้งมากขึ้นจะส่งผลกระทบต่อความเอนเอียงมากขึ้นด้วย เมื่อกลุ่มตัวอย่างมีขนาดเล็ก โมเดลจะมีโอกาสถูกปฏิเสธมากขึ้น และพบว่าการประมาณค่าแบบ SB - χ^2 จะแสดงค่าการทดสอบทางสถิติได้ดีที่สุด

Potthast (1993) ได้ศึกษาวิธีการ CVM (Categorical Variable Method) โดยทำการจำลองข้อมูลที่มีการแจกแจงแบบปกติพหุ (Multivariate Normal Distribution) โดยตัวแปรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับ 5 ระดับ สำหรับโมเดลขนาดต่าง ๆ 4 โมเดล คือ ตั้งแต่ 1 องค์ประกอบ 4 ตัวแปรสังเกตได้ ถึง 4 องค์ประกอบ 16 ตัวแปรสังเกตได้ กลุ่มตัวอย่างมีระดับการแจกแจงที่ไม่เป็นโค้งปกติแตกต่างกัน 4 ระดับ และกลุ่มตัวอย่างมีขนาด 500 และ 1,000 ซึ่งในแต่ละเงื่อนไขของ CVM ใช้ชุดข้อมูล 100 การทดลองซ้ำ (Replications) ผลการศึกษา พบว่า ค่าประมาณของพารามิเตอร์มีความเอนเอียงในระดับน้อย และค่าคลาดเคลื่อนมาตรฐานมีอำนาจการประมาณค่าต่ำ (Underestimated) ในทุก ๆ เงื่อนไข เมื่อขนาดของโมเดลใหญ่ขึ้น หรือตัวแปรมีการแจกแจงแบบโค้งมากขึ้นจะพบว่า มีความเอนเอียงแบบลบ (Negative Biased) มากขึ้น และการทดสอบค่าสถิติ Chi - square มีโอกาสปฏิเสธโมเดลมากขึ้นด้วย การใช้วิธีการ CVM ภายใต้เงื่อนไขโมเดลขนาดใหญ่และกลุ่มตัวอย่างขนาดเล็กนั้น พบว่า ค่าคลาดเคลื่อนมาตรฐานมีอำนาจการประมาณค่าต่ำ (Underestimated) ส่วนค่าสถิติ Chi - square นั้นพบว่า Overestimated

Curran, West and Finch (1996) ได้ใช้การจำลองข้อมูลในการศึกษาเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์ 3 แบบ คือ ML, ADF และ SB - χ^2 ของโมเดล CFA ที่มี 3 องค์ประกอบ

องค์ประกอบ 3 ตัวแปรสังเกตได้ ในเงื่อนไขที่แตกต่างกันดังต่อไปนี้ คือ 1) การกำหนดคุณลักษณะเฉพาะของโมเดลที่ต่างกัน 2 ลักษณะการแจกแจงของพหุตัวแปร (Multivariate Distribution) ตั้งแต่โค้งปกติจนถึงไม่เป็น โค้งปกติในระดับมาก (Severely Non – normal) 3) ขนาดกลุ่มตัวอย่าง 4 ขนาด ตั้งแต่ 100 ถึง 1,000 โดยทำการศึกษาชุดข้อมูลเงื่อนไขละ 200 การทดลองซ้ำ (Replication) ผลการศึกษา พบว่า ค่าสถิติ Chi – square จากวิธีการ $ML - \chi^2$ และ $SB - \chi^2$ มีความเอนเอียงในทุก ๆ ขนาดกลุ่มตัวอย่างภายใต้การแจกแจงพหุตัวแปรเป็น โค้งปกติ และมีค่าเพิ่มสูงขึ้นอย่างมีนัยสำคัญทางสถิติเมื่อระดับการแจกแจงที่ไม่เป็น โค้งปกติเพิ่มมากขึ้น ส่วนค่า $ADF - \chi^2$ จะมีค่าสูงในทุก ๆ เงื่อนไข แม้จะมีการแจกแจงพหุตัวแปรเป็น โค้งปกติก็ตาม ยกเว้น ในกลุ่มตัวอย่างขนาดใหญ่ สำหรับ โมเดลที่ไม่มีการกำหนดคุณลักษณะเฉพาะ (Misspecified) พบว่า ค่า $ML - \chi^2$ จะมีค่าสูงขึ้นเมื่อระดับการแจกแจงที่ไม่เป็น โค้งปกติเพิ่มมากขึ้น แต่ $SB - \chi^2$ และ $ADF - \chi^2$ มีอำนาจการประมาณค่าต่ำ (Underestimated) เมื่อระดับการแจกแจงที่ไม่เป็น โค้งปกติเพิ่มมากขึ้น และสำหรับ โมเดลที่มีการกำหนดคุณลักษณะเฉพาะ (Specified) พบว่า $SB - \chi^2$ จะให้ผลการประมาณค่าที่เหมาะสมกว่าในเงื่อนไขกลุ่มตัวอย่างขนาดกลาง และการแจกแจงตัวแปรที่ไม่เป็น โค้งปกติ

Green, Akey, Fleming, Hershberger and Marquis (1997) ได้ใช้การจำลองข้อมูลในการศึกษาผลกระทบจากจำนวน Categories ของตัวแปรสังเกตได้แบบจำแนกกลุ่ม ในโมเดล CFA ที่มี 1 องค์ประกอบ 20 ตัวแปรสังเกตได้ โดยศึกษาจากประชากรจำนวน 1,000 ชุดข้อมูลที่มีการแจกแจงพหุตัวแปรเป็น โค้งปกติ กลุ่มตัวอย่างถูกสุ่มจากประชากรโดยจำแนกเป็นระดับการแจกแจงที่ไม่เป็น โค้งปกติที่แตกต่างกัน 4 ระดับ จำนวน Categories ต่างกัน คือ 2, 4 และ 6 Categories ผลการศึกษา พบว่า ค่าสถิติ Chi – square มีค่าสูงในทุก ๆ เงื่อนไขของการแจกแจงที่ไม่เป็น โค้งปกติ โดยเฉพาะการแจกแจงที่ไม่เป็น โค้งปกติระดับมากและข้อมูลมีลักษณะเป็น 2 Categories และที่จำนวน 4 และ 6 Categories พบว่า ข้อมูลที่มีการแจกแจงแบบเบ้ซ้ายมากขึ้นจะให้ค่า Chi – square ที่สูงมาก ส่วนข้อจำกัดของการศึกษานี้คือ ใช้ข้อมูลจากประชากรขนาดเล็ก ($N = 1,000$)

Muthén, du Toit and Spisic (1997 cited in Trierweiler, 2009, p. 46) ได้ศึกษาการจำลองข้อมูลโมเดลการวิเคราะห์ระยะยาว (Longitudinal Model) ที่วัดตัวแปรใน 3 ช่วงเวลา สำหรับโมเดล CFA ที่มี 3 องค์ประกอบ 12 ตัวแปรสังเกตได้ ซึ่งมีตัวแปรที่มีลักษณะเป็นแบบ Binary จำนวน 4 ตัวแปร ซึ่งมีตัวแปรอิสระที่ศึกษา ได้แก่ 1) กลุ่มตัวอย่างขนาดต่างกัน 4 ขนาด (200, 400, 800 และ 1,600) 2) การแจกแจงของตัวแปรระดับ Univariate มีลักษณะเบ้ (ความเบ้ตั้งแต่ .08 – .50) และ

3) วิธีการประมาณค่าแบบ Robust Weight Least Square จำนวน 3 แบบ คือ WLSM, WLSMV และ Full WLS ผลการศึกษาพบว่า ในเงื่อนไขกลุ่มตัวอย่างขนาดเล็กและการแจกแจงของตัวแปรแบบเบ้ วิธีการ Full WLS มีความเหมาะสมน้อยกว่าวิธีการ WLSM และ WLSMV ซึ่งในการทดสอบสถิติ Chi-square พบว่า วิธีการ WLSMV ให้ค่าที่ดีกว่าวิธี WLSM และวิธี WLSM จะแสดงอัตราความคลาดเคลื่อนประเภทที่ 1 ที่สูงกว่า

Hutchinson and Olmos (1998) ได้ศึกษาเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์ 2 แบบ คือ ML และ WLS โดยใช้การจำลองข้อมูลประชากรที่ตัวแปรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับ 5 ระดับ และมีการแจกแจงพหุตัวแปรเป็น โค้งปกติ ในเงื่อนไขที่แตกต่างกันดังต่อไปนี้ 1) โมเดลประชากร 2 ลักษณะ คือ โมเดล CFA ที่มี 2 องค์ประกอบ และ 4 องค์ประกอบ (องค์ประกอบละ 4 ตัวแปรสังเกตได้) 2) ขนาดกลุ่มตัวอย่างที่ต่างกัน 2 ขนาด (500 และ 1,000) และ 3) ระดับการแจกแจงตัวแปรของกลุ่มตัวอย่างที่ไม่เป็น โค้งปกติที่ต่างกัน 4 ระดับ โดยวิธีการ ML ใช้เมตริกซ์สหสัมพันธ์แบบเพียร์สัน (PPM) ส่วนวิธีการ WLS ใช้เมตริกซ์สหสัมพันธ์โพลิคอร์ริก (PC) ผลการศึกษา พบว่า วิธีการ ML ให้ผลลัพธ์ความกลมกลืนของโมเดลได้ดีกว่าวิธี WLS ในเงื่อนไขการแจกแจงข้อมูลแบบสมมาตร ส่วนวิธี WLS ให้ผลลัพธ์ความกลมกลืนของโมเดลได้ดีกว่าวิธี ML ในเงื่อนไขการแจกแจงข้อมูลแบบเบ้ ค่าสถิติ Chi-square มีความเอนเอียงในวิธี ML เมื่อข้อมูลมีการแจกแจงแบบเบ้ ในขณะที่วิธี WLS มีความเอนเอียงน้อยมาก ส่วนขนาดของโมเดลและขนาดกลุ่มตัวอย่างส่งผลกระทบต่อค่าสถิติ Chi-square ข้อจำกัดของการศึกษานี้คือ ไม่แสดงผลลัพธ์เกี่ยวกับค่าประมาณของพารามิเตอร์และความคลาดเคลื่อนมาตรฐาน

DiStefano (2002) ได้ศึกษาเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์ 3 แบบ คือ ML, WLS และ Robust ML ในโมเดล CFA ที่ข้อมูลประชากรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับ โดยการจำลองข้อมูลประชากรที่มีการแจกแจงแบบ โค้งปกติโดยประมาณ โดยวิธี ML และ Robust ML ใช้เมตริกซ์สหสัมพันธ์แบบเพียร์สัน (PPM) ส่วนวิธี WLS ใช้เมตริกซ์สหสัมพันธ์โพลิคอร์ริก ซึ่งทำการศึกษาในลักษณะโมเดลที่ต่างกัน 2 แบบ คือ โมเดล 2 องค์ประกอบ 12 ตัวแปรสังเกตได้ (4 และ 8 ตัวแปรในแต่ละองค์ประกอบ), โมเดล 3 องค์ประกอบ 16 ตัวแปรสังเกตได้ (4, 4 และ 8 ตัวแปรในแต่ละองค์ประกอบ ตามลำดับ) ผลการศึกษาพบว่า วิธีการ ML ให้ค่าประมาณของพารามิเตอร์ที่มีความเอนเอียงแบบลบ (Negative Biased) ในทุก ๆ เงื่อนไข เมื่อข้อมูลมีการแจกแจงแบบ โค้งปกติโดยประมาณ ความคลาดเคลื่อนมาตรฐานมีค่าความเอนเอียงสูง โดยเฉพาะเมื่อกำหนดน้ำหนักองค์ประกอบมีค่าปานกลางและกลุ่มตัวอย่างมีขนาดเล็ก ค่าความเอนเอียงของความคลาดเคลื่อนมาตรฐานจะสูงขึ้นในวิธีการ WLS ในเงื่อนไขข้อมูลที่มีการแจกแจงที่ไม่เป็น โค้งปกติ และข้อมูลมีลักษณะจัดกลุ่มแบบเรียงอันดับ การใช้วิธีการ ML โดยใช้เมตริกซ์

สหสัมพันธ์แบบเพียร์สัน (PPM) ในการประมาณค่า นั้น จะให้ค่าประมาณของพารามิเตอร์ ความคลาดเคลื่อนมาตรฐาน และค่าน้ำหนักองค์ประกอบที่มีระดับความเอนเอียงสูงมาก

Oranje (2003 cited in Trierweiler, 2009, p. 48) ได้ศึกษาเปรียบเทียบวิธีการประมาณค่า 5 แบบ ได้แก่ วิธี ML และ WLS ในโปรแกรม LISREL วิธี WLS, WLSM และ WLSMV ในโปรแกรม Mplus โดยการจำลองข้อมูลประชากรที่มีลักษณะจัดกลุ่มแบบเรียงอันดับ ภายใต้เงื่อนไขที่แตกต่างกันของจำนวนตัวแปรแฝงในโมเดล (1, 2 และ 3 ตัวแปร), จำนวนตัวแปร สังเกตได้ต่อ 1 ตัวแปรแฝง (5, 10 และ 15 ตัวแปร), จำนวน Categories ของตัวแปรสังเกตได้ (2 ถึง 5 Categories) และขนาดกลุ่มตัวอย่าง (200, 500 และ 1,000) ซึ่งในแต่ละเงื่อนไขศึกษา กับชุดข้อมูลจำนวน 1,000 การทดลองซ้ำ (Replication) ผลการศึกษาพบว่า ค่า $SB - \chi^2$ และ ค่าความคลาดเคลื่อนมาตรฐาน จากวิธี ML ในโปรแกรม LISREL และวิธี WLSMV แสดงผล ได้เหมาะสมกว่าวิธี WLS หรือ CVM ซึ่ง Oranje ได้สรุปโดยทั่ว ๆ ไปว่า วิธีการ ML นั้น เป็นที่นิยมใช้มากที่สุด และวิธีการ WLSMV นั้น ได้ลดข้อจำกัดเกี่ยวกับกลุ่มตัวอย่างขนาดเล็ก และโมเดลขนาดใหญ่ได้ดี

Flora and Curran (2004, p. 466) ได้ทำการศึกษาผลจากการประมาณค่าพารามิเตอร์ด้วย เทคนิคการประมาณค่าแบบ WLS และ Robust WLS สำหรับโมเดล CFA ที่ตัวแปรมีมาตรวัดระดับ เรียงอันดับ โดยวิธีการจำลองข้อมูล และกำหนดเงื่อนไขที่แตกต่างกัน คือ กลุ่มตัวอย่าง 4 ขนาด (100, 200, 500 และ 1,000) ระดับการแจกแจงข้อมูล 5 ระดับ การกำหนดลักษณะเฉพาะของโมเดล (Model Specification) 4 แบบ และจำนวน Categories ของตัวแปร 2 แบบ ซึ่งทำการทดลองซ้ำ (Replicate) 500 ครั้ง ในแต่ละเงื่อนไข โดยใช้โปรแกรม EQS ผลการศึกษาพบว่า การใช้เมตริกซ์ สหสัมพันธ์โพลีคอร์ริก (Polychoric Correlation) มีความแข็งแกร่งต่อการละเมิดข้อตกลงเบื้องต้น เกี่ยวกับการแจกแจงแบบปกติ วิธีการประมาณค่าพารามิเตอร์แบบ WLS แสดงผลลัพธ์ที่เหมาะสม เฉพาะในกรณีที่กลุ่มตัวอย่างมีขนาดใหญ่ ($n = 1,000$) เท่านั้น ส่วนเมื่อกลุ่มตัวอย่างมีขนาดเล็ก จะมีประสิทธิภาพของการประมาณค่าต่ำ ส่วนวิธีการ Robust WLS แสดงผลลัพธ์ที่เหมาะสม ในทุก ๆ เงื่อนไข Flora & Curran ได้สรุปและเสนอแนะว่าวิธีการ Robust WLS มีความเหมาะสม ในการประมาณค่าในเงื่อนไขที่ตัวแปรมีมาตรวัดแบบเรียงอันดับ ในโมเดลขนาดกลาง ถึงขนาดใหญ่ และกลุ่มตัวอย่างขนาดกลางถึงขนาดเล็ก

Beauducel and Herzberg (2006) ได้ศึกษาเปรียบเทียบวิธีการประมาณค่า 2 แบบ คือ ML และ WLSMV (Robust Weight Least Squares with Means and Variance Adjusted) ในโมเดล CFA ที่ข้อมูลประชากรมีการแจกแจงเป็นโค้งปกติ โดยตัวแปรอิสระที่ศึกษา คือ 1) จำนวนองค์ประกอบ ที่แตกต่างกัน ตั้งแต่ 1 ถึง 8 องค์ประกอบ ซึ่งมีตัวแปรสังเกตได้ 5 ตัวแปร ต่อองค์ประกอบ

2) กลุ่มตัวอย่างขนาดต่าง ๆ (ตั้งแต่ 250 จนถึง 1,000) 3) จำนวน Categories (2 ถึง 6 Categories) ส่วนตัวแปรตามที่ศึกษา ได้แก่ ค่าประมาณของพารามิเตอร์, ค่าความคลาดเคลื่อนมาตรฐานและดัชนีวัดความกลมกลืนต่าง ๆ ผลการศึกษา พบว่า สำหรับ โมเดลที่ข้อมูลจัดอันดับมีจำนวน 2 และ 3 Categories วิธีการ WLSMV มีอัตราการใช้ปฏิเสธที่สอดคล้องกับอัตราการใช้ปฏิเสธในค่าคาดหวัง (Expected Rejection Rate) มากกว่าวิธีการ ML ที่ระดับนัยสำคัญ .05 และวิธีการ WLSMV ให้ค่าน้ำหนักองค์ประกอบที่สูง ค่าประมาณพารามิเตอร์มีความแม่นยำมากกว่าผลการศึกษา ยังพบอีกว่าวิธีการประมาณค่าและจำนวน Categories นั้นมีผลต่อค่าดัชนี CFI ซึ่งค่านี้จะเพิ่มขึ้นที่เงื่อนไขโมเดลมีจำนวน 2 และ 3 Categories เมื่อจำนวน Categories ที่มาก ๆ พบว่าวิธีการประมาณค่าทั้งสองแบบให้ผลใกล้เคียงกัน ในเงื่อนไขจำนวน 2 และ 3 Categories วิธีการ WLSMV ให้ค่าดัชนี TLI ที่ดีกว่า แต่ดัชนี TLI จะมีค่าเพิ่มขึ้นเมื่อจำนวน Categories และขนาดกลุ่มตัวอย่างเพิ่มขึ้นในวิธีการ ML เท่านั้น โดยภาพรวม Beauducel & Herzberg ให้ข้อสรุปว่า ในเงื่อนไขที่ข้อมูลมีลักษณะจัดกลุ่มแบบเรียงอันดับ จำนวน 2 หรือ 3 Categories วิธีการ WLSMV ให้ผลการประมาณค่าที่ดีกว่าวิธีการ ML ในผลลัพธ์ค่าประมาณของพารามิเตอร์, Rate of Non-convergence และค่าดัชนีวัดความกลมกลืน เช่น CFI, TLI และ RMSEA ซึ่งข้อจำกัดของการศึกษาครั้งนี้คือ การทดสอบโมเดลมีเงื่อนไขภายใต้การแจกแจงเป็นโค้งปกติ ไม่ได้ศึกษาผลกระทบจากการแจกแจงข้อมูลแบบเบ้

Trierweiler (2009, pp. 136–147) ได้ศึกษาประสิทธิภาพการประมาณค่าพารามิเตอร์โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันสำหรับข้อมูลจัดกลุ่มแบบเรียงอันดับในโปรแกรม LISREL โดยการจำลองข้อมูล โดยเปรียบเทียบเงื่อนไขที่แตกต่างกัน คือ 1) วิธีการประมาณค่าพารามิเตอร์ 4 แบบ คือ ML, Robust ML, WLS และ Robust DWLS 2) โมเดลประชากรที่แตกต่างกัน 2 โมเดล คือ โมเดล 3 องค์ประกอบ 9 ตัวแปรสังเกตได้ และ 3 องค์ประกอบ 18 ตัวแปรสังเกตได้ 3) ระดับการแจกแจงของข้อมูล 5 ระดับ และ 4) ขนาดกลุ่มตัวอย่าง 4 ขนาด ($n = 100, 200, 500$ และ $1,000$) โดยทำการศึกษาเงื่อนไขละ 500 การทดลองซ้ำ (Replications) และตัวแปรตามที่ทำการศึกษา ได้แก่ 1) ความเอนเอียงของค่าประมาณของพารามิเตอร์ (Parameter Estimates Bias) 2) ความเอนเอียงของค่าคลาดเคลื่อนมาตรฐาน (Standard Error Bias) 3) ค่าสถิติ Chi-square และ 4) ดัชนีวัดความสอดคล้องกลมกลืน ผลการศึกษา พบว่า

1. วิธีการ ML (Normal Theory ML) ในเงื่อนไขที่ข้อมูลมีการแจกแจงแบบปกติพหุ แสดงค่าประมาณของพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐานที่มีความเอนเอียงเล็กน้อย ยกเว้นในกลุ่มตัวอย่างขนาดเล็ก ($n=100$) ที่พบว่า ค่าเบี่ยงเบนมาตรฐาน แสดงค่า Underestimate ในระดับสูง เมื่อระดับการแจกแจงที่ไม่เป็นโค้งปกติมากขึ้นจะแสดงค่าความเอนเอียงสูงขึ้นด้วย

โดยเฉพาะในกลุ่มตัวอย่างขนาดเล็ก ขนาดของโมเดลและขนาดกลุ่มตัวอย่างส่งผลต่อค่าความเอนเอียงอย่างมีปฏิสัมพันธ์กัน คือ เมื่อกลุ่มตัวอย่างเพิ่มขึ้นค่าความเอนเอียงจะลดลง และเมื่อจำนวนตัวแปรสังเกตได้ต้องประกอบมากขึ้นค่าความเอนเอียงของค่าคลาดเคลื่อนมาตรฐานจะลดลง ค่าสถิติ Chi – square และดัชนีวัดความสอดคล้องกลมกลืน คือ CFI และ RMSEA แสดงค่าไม่เหมาะสมในทุก ๆ เงื่อนไข โดยสรุป ในเงื่อนไขข้อมูลจัดกลุ่มแบบเรียงอันดับที่มีการแจกแจงที่ไม่เป็นโค้งปกติ วิธีการ ML ให้ผลการประมาณค่าที่มีประสิทธิภาพต่ำ

2. วิธีการ WLS (Weighted Least Square) ไม่พบความเอนเอียงสำหรับค่าประมาณพารามิเตอร์ ค่าคลาดเคลื่อนมาตรฐาน และค่าสถิติ Chi – square ในกลุ่มตัวอย่างขนาดใหญ่ ($n = 1000$) เท่านั้น เมื่อขนาดกลุ่มตัวอย่างลดลง และขนาดของ โมเดลใหญ่ขึ้นจะพบความเอนเอียงสูงขึ้นอย่างมีนัยสำคัญ ดัชนีวัดความสอดคล้องกลมกลืนแสดงค่าไม่เหมาะสมในกลุ่มตัวอย่างขนาดเล็ก โมเดลขนาดใหญ่และที่ระดับการแจกแจงที่ไม่เป็นโค้งปกติมากขึ้น โดยสรุป ในเงื่อนไขข้อมูลจัดกลุ่มแบบเรียงอันดับที่มีการแจกแจงที่ไม่เป็นโค้งปกติ วิธีการ WLS ให้ผลการประมาณค่าที่มีประสิทธิภาพต่ำในกลุ่มตัวอย่างขนาดเล็ก ($n < 500$)

3. วิธีการ Robust ML แสดงค่าประมาณของพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐานที่มีความเอนเอียงเล็กน้อย (ต่ำกว่า 5%) ยกเว้นในกลุ่มตัวอย่างขนาดเล็ก ($n = 100$) ระดับการแจกแจงที่ไม่เป็นโค้งปกติและจำนวนตัวแปรสังเกตได้ต้องประกอบหรือขนาดของโมเดลที่มากขึ้น ส่งผลกระทบต่อความเอนเอียงให้มีค่าสูงขึ้น ส่วนดัชนีวัดความสอดคล้องกลมกลืนแสดงค่าเหมาะสมในทุกเงื่อนไข โดยสรุป วิธีการ Robust ML ให้ผลการประมาณค่าที่มีประสิทธิภาพดี ยกเว้น ในเงื่อนไขกลุ่มตัวอย่างขนาดเล็ก การแจกแจงที่ไม่เป็นโค้งปกติระดับมากและโมเดลขนาดใหญ่

4. วิธีการ Robust DWLS แสดงค่าประมาณของพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐานที่มีความเอนเอียงเล็กน้อย (ต่ำกว่า 5%) ในทุกเงื่อนไข ขนาดของโมเดลและขนาดกลุ่มตัวอย่างส่งผลต่อประสิทธิภาพเพียงเล็กน้อย ส่วนดัชนีวัดความสอดคล้องกลมกลืนแสดงค่าเหมาะสมในทุกเงื่อนไข โดยสรุป วิธีการ Robust DWLS ให้ผลการประมาณค่าที่มีประสิทธิภาพดีมีความแกร่งต่อเงื่อนไขกลุ่มตัวอย่างขนาดเล็ก การแจกแจงที่ไม่เป็นโค้งปกติและโมเดลขนาดใหญ่

Lei (2009) ได้ศึกษาเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์ 2 แบบ คือ Robust ML และ WLSMV (Robust Weight Least Squares with Means and Variance Adjusted) สำหรับโมเดล CFA ใน 2 ลักษณะคือ กำหนดลักษณะเฉพาะ (Specified) และไม่กำหนดลักษณะเฉพาะ (Misspecified) โดยข้อมูลประชากรมีลักษณะเป็นข้อมูลจัดกลุ่มแบบเรียงอันดับและมีการแจกแจง

เป็นโค้งปกติโดยประมาณ กลุ่มตัวอย่างถูกสุ่มจากประชากรโดยจำแนกการแจกแจงแบบเบ้ในระดับต่าง ๆ และกำหนดขนาดของโมเดล 2 ลักษณะ คือ โมเดล 2 องค์ประกอบ 6 ตัวแปรสังเกตได้ และโมเดล 3 องค์ประกอบ 9 ตัวแปรสังเกตได้ โดยตัวแปรตามที่ศึกษา ได้แก่ Type I Error Rate ของสถิติ Chi – square, Non – convergence Rate ของโมเดล, Relative Bias ของค่าประมาณพารามิเตอร์และความคลาดเคลื่อนมาตรฐานของโมเดล ผลการศึกษา พบว่า เมื่อขนาดของโมเดลลดลงและความเบ้ของการแจกแจงข้อมูลเพิ่มขึ้น วิธีการ WLSMV จะแสดงค่าร้อยละของ Improper Solutions มากขึ้น ร้อยละของการปฏิเสธโมเดลมีค่าน้อย และค่าความเอนเอียงทางลบของความคลาดเคลื่อนมาตรฐานมากกว่าวิธีการ Robust ML อย่างไรก็ตามวิธีการ WLSMV ก็ให้ผลที่ดีกว่าวิธีการ ML ในการควบคุม Type I error ในภาพรวมและแกร่งต่อโมเดลที่ไม่กำหนดลักษณะเฉพาะ ซึ่งวิธีการ WLSMV จะมีความอ่อนไหวต่อกลุ่มตัวอย่างขนาดเล็กและระดับการแจกแจงแบบเบ้มากกว่าวิธีการ Robust ML นอกจากนี้สำหรับโมเดลที่มีการกำหนดลักษณะเฉพาะอย่างถูกต้อง พบว่า ไม่มีความเอนเอียงในค่าประมาณของพารามิเตอร์ในวิธีการประมาณค่าทั้งสองแบบ สำหรับกลุ่มตัวอย่างขนาดเล็ก ($n = 100$) พบว่า วิธีการ Robust ML แสดงผลการประมาณค่าความคลาดเคลื่อนมาตรฐานได้ดีกว่าวิธีการ WLSMV เล็กน้อย ความแตกต่างของความเอนเอียงของความคลาดเคลื่อนมาตรฐานมีค่าน้อยมากเมื่อข้อมูลโมเดลมีการแจกแจงเป็นโค้งปกติ วิธีการ WLSMV ที่การแจกแจงมีความเบ้มากขึ้นในโมเดลขนาดเล็กและกลุ่มตัวอย่างขนาดเล็ก ($n = 100$) ความเอนเอียงของความคลาดเคลื่อนมาตรฐาน และค่า Type I Error มีค่าสูงขึ้นมาก

Muthén and Kaplan (1985); Potthast (1993); DiStafano (2002); Flora and Curran (2004) กล่าวว่า การใช้การจำลองข้อมูลในการศึกษาประสิทธิภาพของวิธีการประมาณค่าพารามิเตอร์ตามทฤษฎีที่พัฒนาขึ้นเพื่อดำเนินการกับข้อมูลที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) เช่น วิธีการ CVM, WLS หรือ Robust WLS (WLSM, WLSMV) เมื่อการจำลองข้อมูลใช้กลุ่มตัวอย่างขนาดใหญ่ การแจกแจงข้อมูลเป็นโค้งปกติและมีจำนวน Categories ตั้งแต่ 5 ขึ้นไป จะพบว่า ค่าประมาณของพารามิเตอร์มีอำนาจการประมาณค่าต่ำเพียงเล็กน้อย โดยทั่วไปเมื่อการแจกแจงข้อมูลไม่เป็นโค้งปกติ วิธีการ CVM หรือ Robust WLS จะให้ค่าประมาณของพารามิเตอร์ที่มีความแกร่งในกรณีที่โมเดลมีขนาดเล็ก (ตัวแปรสังเกตได้มีจำนวนน้อย) และกลุ่มตัวอย่างมีขนาดเล็ก ค่าสถิติ Chi – square มีค่าสูงเมื่อระดับการแจกแจงมีความเบ้มากขึ้น ในกรณีที่โมเดลมีขนาดใหญ่และกลุ่มตัวอย่างมีขนาดเล็ก จะให้ค่าความคลาดเคลื่อนมาตรฐานที่มีความเอนเอียงทางลบมากขึ้น

Jöreskog (1990) และ Muthén (1983; 1984) กล่าวว่า ถ้าข้อมูลมีการแจกแจงที่ไม่เป็นโค้งปกติ และกลุ่มตัวอย่างมีขนาดใหญ่ วิธีการ WLS หรือ CVM เป็นวิธีการที่มีความน่าเชื่อถือสำหรับข้อมูลที่มีลักษณะจำแนกกลุ่ม แต่ต้องขึ้นอยู่กับขนาดของกลุ่มตัวอย่างที่ต้องมีขนาดใหญ่ และจำนวนตัวแปรสังเกตได้ใน โมเดล ในกรณีที่กลุ่มตัวอย่างมีขนาดเล็กถึงปานกลางและโมเดลมีขนาดใหญ่ วิธีการ Robust WLS นั้นจะมีความเหมาะสมกว่า อย่างไรก็ตามวิธีการนี้ยังเป็นวิธีการที่ใหม่ มีการศึกษาวิจัยเกี่ยวกับผลกระทบของวิธีการนี้จำนวนจำกัด

MacKinnon et al. (2004 cited in Muthén, 2010, p. 6 – 12) ศึกษาเปรียบเทียบวิธีการประมาณค่าแบบ ML และ Bayesian ในโมเดล SEM ที่ข้อมูลมีการแจกแจงไม่เป็นโค้งปกติ พบว่าวิธีการ ML ในเงื่อนไขละเมิดข้อตกลงเบื้องต้นในการแจกแจงข้อมูลเป็นโค้งปกติแสดงผลลัพธ์ค่าประมาณของพารามิเตอร์ที่ไม่ถูกต้อง ในขณะที่วิธี Bayesian ทำการวิเคราะห์ข้อมูลโดยการยอมให้การแจกแจงภายหลัง (Posterior Distribution) สามารถมีการแจกแจงไม่เป็นโค้งปกติได้ แสดงผลลัพธ์ค่าประมาณของพารามิเตอร์ที่ดีกว่า

Asparouhov and Muthén (2010, pp. 2 – 5) กล่าวถึงการวิเคราะห์องค์ประกอบสำหรับตัวชี้วัดแบบต่อเนื่อง (Factor Analysis with Continuous Indicators) ว่าวิธีการประมาณค่าแบบ Bayesian ในโมเดล CFA ที่มี 1 องค์ประกอบ โดยเปรียบเทียบโมเดล 2 ขนาด คือ 30 ตัวแปรสังเกตได้ และ 5 ตัวแปรสังเกตได้ และเปรียบเทียบลักษณะการกำหนดการประมาณค่าพารามิเตอร์ (Parameterization) ใน 3 ลักษณะ คือ ลักษณะที่ 1 “L” หมายถึง ค่าน้ำหนักองค์ประกอบทุกค่าถูกประมาณค่าและค่าความแปรปรวนขององค์ประกอบกำหนดให้เท่ากับ 1.00 ลักษณะที่ 2 “V” หมายถึง กำหนดค่าน้ำหนักองค์ประกอบค่าแรกให้มีค่าเท่ากับ 1.00 และทำการประมาณค่าความแปรปรวนขององค์ประกอบทุกค่า และลักษณะที่ 3 “PX” หมายถึง ค่าน้ำหนักองค์ประกอบและความแปรปรวนขององค์ประกอบทุกค่าถูกประมาณค่าไม่มีการกำหนดค่าใด ๆ รวมถึงศึกษาในขนาดกลุ่มตัวอย่างที่แตกต่างกัน 6 ขนาด คือ 50 100 200 500 1,000 และ 5,000 ผลการศึกษาพบว่า ลักษณะ “PX” มีประสิทธิภาพที่ดีกว่าลักษณะอื่น ๆ ในกลุ่มตัวอย่างขนาดเล็ก ($n = 50$ และ $n = 100$) และ ลักษณะ “PX” ไม่เกิดค่าความเอนเอียงในทุก ๆ ขนาดกลุ่มตัวอย่าง ลักษณะ “L” พบค่า ความเอนเอียงในกลุ่มตัวอย่างขนาดเล็กแต่ในกลุ่มตัวอย่าง $n = 200$ ขึ้นไป จะให้ผลลัพธ์ที่ดี ลักษณะ “V” แสดงผลลัพธ์ที่แย่กว่าลักษณะอื่น ๆ ในเงื่อนไขจำนวนตัวแปรสังเกตได้ในโมเดล พบว่า จำนวนตัวแปรสังเกตได้ทีมากจะแสดงค่าความเอนเอียงสูงกว่าและมีประสิทธิภาพต่ำกว่า จำนวนตัวแปรสังเกตได้น้อย ๆ แต่เมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้นค่าความเอนเอียงจะลดลงและผลลัพธ์แสดงผลดีขึ้น

Asparouhov and Muthén (2010, pp. 18 – 22) กล่าวถึงการเปรียบเทียบวิธีการประมาณค่าแบบ WLSMV และแบบ Bayesian ในโมเดล SEM ที่ตัวแปรมีลักษณะจำแนกกลุ่ม (Categorical Variable) และข้อมูลขาดหาย (Missing Data) ผลการศึกษา พบว่า วิธี WLSMV แสดงค่าความเอนเอียงและมีร้อยละของ Coverage ต่ำ ในขณะที่วิธี Bayesian ไม่พบความเอนเอียงและร้อยละของ Coverage มีค่าเฉลี่ยประมาณ 95% ซึ่งการประมาณค่าเข้าใกล้ค่าจริงของพารามิเตอร์มากกว่า

Asparouhov and Muthén (2010, pp. 45 – 47) กล่าวถึงการเปรียบเทียบประสิทธิภาพวิธีการประมาณค่าแบบ Bayesian และ WLSMV ใน Multiple Indicator Growth Modeling สำหรับตัวแปรจำแนกกลุ่ม โดยจำนวนพารามิเตอร์ที่ทำการประมาณค่า 36 พารามิเตอร์ ผลการศึกษา พบว่าวิธีการประมาณค่าทั้งสองวิธี ไม่พบความเอนเอียงและช่วงเชื่อมั่นของ Coverage ประมาณ 95% วิธีการ Bayesian แสดงค่า Mean Square Error ที่ต่ำกว่าวิธีการ WLSMV ในทุก ๆ ค่าประมาณพารามิเตอร์ ซึ่งสามารถสรุปได้ว่า วิธีการประมาณค่าพารามิเตอร์แบบ Bayesian มีความแม่นยำมากกว่าวิธีการ WLSMV และใช้เวลาในการคำนวณน้อยกว่าประมาณ 1.5 เท่า

Chumney (2012) ได้ศึกษาการเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์แบบ Bayesian, Partial Least Square (PLS) และ Generalized Structured Component Analysis (GSCA) กับวิธี Maximum Likelihood (ML) สำหรับโมเดลสมการเชิงโครงสร้างในเงื่อนไขกลุ่มตัวอย่างขนาดเล็ก ผลการศึกษา พบว่า ในขนาดกลุ่มตัวอย่าง $n = 360$ ค่าประมาณของพารามิเตอร์จากวิธี Bayesian ไม่พบค่าความเอนเอียงหรือไม่แตกต่างจากวิธีการ ML ในขณะที่วิธีการ PLS และ GSCA แตกต่างจากวิธีการ ML ด้วยค่าความเอนเอียงสมพัทธ์ 17.90% ภายใต้กลุ่มตัวอย่างขนาดเล็ก $n < 100$ พบว่า ทุกวิธีการประมาณค่ามีความเอนเอียงจากวิธีการ ML อย่างมีนัยสำคัญทางสถิติ ส่วนค่าดัชนีวัดความสอดคล้องกลมกลืนของโมเดล (Fit Index) พบว่า ทุกวิธีการประมาณค่าพารามิเตอร์แสดงค่าเหมาะสมดี ไม่แตกต่างกัน