

บทที่ 5

สรุปผล อภิปรายผล และข้อเสนอแนะ

การวิจัยเรื่อง ประสิทธิภาพของวิธีการประมาณค่าพารามิเตอร์แบบ ML, WLSMV และ Bayesian สำหรับข้อมูลจำแนกกลุ่มแบบเรียงอันดับ ที่มีการแจกแจงไม่เป็นโค้งปกติ ภายใต้เงื่อนไขที่แตกต่างกัน: กรณีการวิเคราะห์หองค์ประกอบเชิงยืนยันอันดับแรก ผู้วิจัยนำเสนอสาระสำคัญตามลำดับ คือ วัตถุประสงค์ในการวิจัย สมมติฐานการวิจัย วิธีดำเนินการวิจัย สรุปผลการวิจัย อภิปรายผลการวิจัย ข้อเสนอแนะในการนำผลการวิจัยไปใช้ และข้อเสนอแนะสำหรับการวิจัยครั้งต่อไป ซึ่งมีรายละเอียดดังต่อไปนี้

วัตถุประสงค์ของการวิจัย

เพื่อศึกษาและเปรียบเทียบประสิทธิภาพของวิธีประมาณค่าพารามิเตอร์แบบ ML ในกรณีที่ไม่มี การแปลงข้อมูลและกรณีที่มีการแปลงข้อมูลให้มีการแจกแจงเป็นโค้งปกติ กับวิธีการประมาณค่าพารามิเตอร์แบบ WLSMV และวิธีการ Bayesian ในกรณีที่ข้อมูลการวัดอยู่ในมาตรวัดจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกติ ภายใต้เงื่อนไขของขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด ได้แก่ ขนาด 80, 160, 370 และ 623 หน่วยตัวอย่าง

สมมติฐานของการวิจัย

1. ภายใต้เงื่อนไขขนาดกลุ่มตัวอย่างเท่ากัน เมื่อใช้วิธีการประมาณค่าพารามิเตอร์ต่างกัน จะให้ผลลัพธ์ของการวิเคราะห์ที่มีความเที่ยงตรงและความเอนเอียงแตกต่างกัน โดยวิธี WLSMV และวิธี Bayesian จะให้ผลลัพธ์ของการวิเคราะห์ที่มีความเที่ยงตรงสูงกว่า/ ความเอนเอียงที่ต่ำกว่าวิธีการ ML ทั้งสองกรณี
2. เมื่อใช้วิธีการประมาณค่าพารามิเตอร์เดียวกัน ภายใต้เงื่อนไขขนาดกลุ่มตัวอย่างที่แตกต่างกัน จะให้ผลลัพธ์ของการวิเคราะห์ที่มีความเที่ยงตรงและความเอนเอียงแตกต่างกัน โดยวิธีการประมาณค่า ทั้ง 4 วิธี จะให้ผลลัพธ์ของการวิเคราะห์ที่มีความเที่ยงตรงสูงขึ้นและความเอนเอียงน้อยลงเมื่อใช้กับกลุ่มตัวอย่างขนาดใหญ่ และผลลัพธ์ของการวิเคราะห์จะมีความเที่ยงตรงน้อยลงและความเอนเอียงเพิ่มสูงขึ้นเมื่อใช้กับกลุ่มตัวอย่างขนาดเล็ก

วิธีดำเนินการวิจัย

1. ข้อมูลประชากรที่ใช้ในการวิจัยในครั้งนี้ เป็นข้อมูลทุติยภูมิ (Secondary Source) โดยอ้างอิงจากงานวิจัยเชิงสำรวจ โดยผู้วิจัยคัดเลือกชุดข้อมูลที่มีความสมบูรณ์ใช้เป็นข้อมูลประชากรเทียมจำนวน 10,000 ชุดข้อมูล ที่มีลักษณะข้อมูลจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) ที่ใช้มาตรวัด Likert Type Scale 5 ระดับ
2. ทดสอบการแจกแจงปกติของแต่ละตัวแปร (Univariate Normality Test) และการทดสอบการแจกแจงปกติพหุ (Multivariate Normality Test) ของข้อมูลประชากร ซึ่งพบว่ามีการแจกแจงไม่เป็นโค้งปกติ (Non – normality Distribution) คือ เบ้ทางซ้าย (Negative Skewness) ทุกตัวแปร
3. โมเดลที่ใช้เป็นต้นแบบในการศึกษา ผู้วิจัยดำเนินการสร้างโมเดลการวัด หรือโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันอันดับแรก (First Order Confirmatory Factor Analysis) ที่ประกอบด้วย 2 องค์ประกอบ องค์ประกอบละ 4 ตัวแปรสังเกตได้ รวมทั้งสิ้น 8 ตัวแปรสังเกตได้
4. คำนวณค่าพารามิเตอร์ด้วยวิธีการ WLSMV โดยกำหนดให้เป็น โมเดลพารามิเตอร์ที่ 1 เพื่อใช้เปรียบเทียบการประมาณค่าพารามิเตอร์จากกลุ่มตัวอย่างด้วยวิธีการ ML, ML_{Tr} และ WLSMV และคำนวณค่าพารามิเตอร์ด้วยวิธีการ Bayesian โดยกำหนดให้เป็น โมเดลพารามิเตอร์ที่ 2 เพื่อใช้เปรียบเทียบการประมาณค่าพารามิเตอร์จากกลุ่มตัวอย่างด้วยวิธีการ Bayesian
5. สุ่มข้อมูลจากประชากร ($N = 10,000$) ภายใต้งैอนไขกลุ่มตัวอย่างที่ต่างกัน 4 ขนาด คือ 80, 160, 370 และ 623 หน่วยตัวอย่าง เงื่อนไขละ 200 การทดลองซ้ำ (Replications)
6. ดำเนินการประมาณค่าพารามิเตอร์ในแต่ละเงื่อนไข โดยการกำหนด โมเดล และการปรับ โมเดล จะดำเนินการเหมือนกับ โมเดลพารามิเตอร์ทุกประการ ซึ่งวิธีการประมาณค่าพารามิเตอร์ 4 วิธี ประกอบด้วย
 - 6.1 ML เป็นวิธีการประมาณค่าพารามิเตอร์ภายใต้การละเมิดข้อตกลงเบื่องต้นเกี่ยวกับการแจกแจงแบบโค้งปกติ และข้อมูลแบบต่อเนื่อง
 - 6.2 ML_{Tr} เป็นวิธีการเปลี่ยนรูปข้อมูล (Data Transforming) ในแต่ละตัวแปร โดยวิธีการคำนวณ Normal Score ด้วยโปรแกรม LISREL Version 8.7 ซึ่งทำให้ข้อมูลมีการแจกแจงเป็นโค้งปกติและมีลักษณะเป็นข้อมูลแบบต่อเนื่องตามข้อตกลงเบื่องต้นเสียก่อนแล้วดำเนินการประมาณค่าพารามิเตอร์ด้วยวิธีการ ML
 - 6.3 WLSMV เป็นวิธีการประมาณค่าพารามิเตอร์ที่มีความแกร่งต่อการละเมิดข้อตกลงเบื่องต้นเกี่ยวกับลักษณะข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกติ และกลุ่มตัวอย่างขนาดเล็ก

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

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

7.1 ค่าความเอนเอียงของค่าประมาณพารามิเตอร์ (Parameter Estimates Bias)

7.2 ค่าความเอนเอียงของค่าคลาดเคลื่อนมาตรฐาน (Standard Error Bias)

7.3 ค่าความเชื่อมั่น (Reliability) ได้แก่ ค่า ρ_c และค่า ρ_v

7.4 ร้อยละความสอดคล้องของโมเดล จากดัชนีวัดความสอดคล้องกลมกลืนของโมเดล (Model Fit Indices)

สรุปผลการวิจัย

ผลการวิเคราะห์ข้อมูลพิจารณาตามวิธีการประมาณค่าพารามิเตอร์ที่ต่างกัน คือ ML, ML_h, WLSMV และ Bayesian สามารถแสดงผลการวิเคราะห์ได้ดังนี้

1. วิธีการประมาณค่าพารามิเตอร์แบบ ML (Maximum Likelihood)

1.1 ค่าทำนายของค่าประกอบจากการประมาณค่ามีค่าเฉลี่ยต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) และช่วงความเชื่อมั่น 95 % ไม่ครอบคลุมค่าพารามิเตอร์ในทุก ๆ ขนาดกลุ่มตัวอย่าง ส่วนค่าคลาดเคลื่อนมาตรฐานของค่าเฉลี่ยจะมีค่าลดลงเมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้น

1.2 ความแตกต่างระหว่างค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์

1.2.1 $n = 80$ พบว่า ในแต่ละค่าทำนายของค่าประกอบ มีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกค่าทำนายของค่าประกอบ

1.2.2 $n = 160$ พบว่า ในแต่ละค่าน้ำหนักองค์ประกอบมีความแตกต่างกัน อย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกค่าน้ำหนักองค์ประกอบ

1.2.3 $n = 370$ พบว่า ในแต่ละค่าน้ำหนักองค์ประกอบมีความแตกต่างกัน อย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกค่าน้ำหนักองค์ประกอบ

1.2.4 $n = 623$ พบว่า ในแต่ละค่าน้ำหนักองค์ประกอบมีความแตกต่างกัน อย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกค่าน้ำหนักองค์ประกอบ

สรุป ความแตกต่างระหว่างค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ จากวิธีการ ML พบว่า ค่าประมาณของพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์อย่างมีนัยสำคัญทางสถิติ ($p < .001$) ในทุกขนาดกลุ่มตัวอย่าง แม้ว่าจะเข้าใกล้ค่าจริงของพารามิเตอร์มากขึ้น เมื่อกลุ่มตัวอย่างมากขึ้นก็ตาม

1.3 ค่าร้อยละความเอนเอียงสัมพัทธ์ (Relative Bias) ของค่าประมาณพารามิเตอร์

1.3.1 $n = 80$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 11.010 – 15.598 ส่วนใหญ่มีค่าอยู่ในระดับสูง มีอัตราส่วนของระดับสูง: ปานกลาง: ต่ำ = 5: 4: 0

1.3.2 $n = 160$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 5.310 – 11.859 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 3: 6: 0

1.3.3 $n = 370$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 4.586 – 11.007 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 2: 6: 1

1.3.4 $n = 623$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 4.781 – 9.855 ซึ่งมีค่าอยู่ในระดับปานกลางทุกค่าน้ำหนักองค์ประกอบ ยกเว้น ค่าน้ำหนักองค์ประกอบ A_{B2} ที่มีค่าอยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 8: 1

สรุป ในวิธีการ ML ส่วนใหญ่ค่าความเอนเอียงสัมพัทธ์มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นในทุกค่าน้ำหนักองค์ประกอบ ยกเว้น ค่าน้ำหนักองค์ประกอบ A_{A3} A_{A4} และ A_{B1} ที่ขนาดตัวอย่าง 370 ที่ค่าความเอนเอียงสัมพัทธ์ต่ำกว่าขนาดตัวอย่าง 623

1.4 ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) ของค่าประมาณพารามิเตอร์

1.4.1 $n = 80$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 11.010 – 15.598 มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับสูงทุกค่าน้ำหนักองค์ประกอบ มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 9: 0: 0

1.4.2 $n = 160$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 7.527 – 10.728 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 3: 6: 0

1.4.3 $n = 370$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 4.711 – 6.827 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 7: 2

1.4.4 $n = 623$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 4.411 – 5.040 ซึ่งมีค่าอยู่ในระดับต่ำทุกค่า ยกเว้น น้ำหนักองค์ประกอบ A_{43} มีค่าอยู่ในระดับปานกลาง อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 1: 8

สรุปในวิธีการ ML พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นทุกค่าน้ำหนักองค์ประกอบ และที่ขนาดกลุ่มตัวอย่างเท่ากับ 80 จะให้ค่าที่ไม่เหมาะสม

1.5 ค่าความเชื่อมั่น (Reliability) ของโมเดลการวัด ได้แก่ ค่าความเชื่อมั่นของตัวแปรแฝง (Construct Reliability: ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (Average Variance Extracted: ρ_v) ของโมเดลการวัด

1.5.1 $n = 80$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

1.5.2 $n = 160$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

1.5.3 $n = 370$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

1.5.4 $n = 623$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

สรุป ในวิธีการ ML พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001 ในทุกขนาดกลุ่มตัวอย่าง แม้ว่าเมื่อกลุ่มตัวอย่างเพิ่มมากขึ้น จะให้ค่าความเชื่อมั่นของโมเดลการวัดที่เข้าใกล้ค่าจริงของพารามิเตอร์ก็ตาม

1.6 ร้อยละของ โมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit)

เมื่อพิจารณาจากค่าดัชนีวัดความสอดคล้องกลมกลืน (CFI) วิธีการ ML มีอัตราร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ ในขนาดกลุ่มตัวอย่าง 80, 160, 370 และ 623 พบว่ามีค่า 36.0%, 48.0%, 61.0% และ 75.5% เพิ่มขึ้นตามลำดับ และเมื่อพิจารณาจากค่าดัชนี TLI พบว่ามีค่า 48.0%, 60.5%, 73.0% และ 81.0% เพิ่มขึ้นตามลำดับ แต่อย่างไรก็ตาม แม้ว่าร้อยละของโมเดลที่สอดคล้องจะเพิ่มขึ้นตามขนาดของกลุ่มตัวอย่างที่เพิ่มขึ้น แต่ในภาพรวมแล้ว พบว่า ร้อยละของโมเดลที่ไม่สอดคล้องก็ยังคงมีอยู่เป็นจำนวนมาก

2. วิธีการประมาณค่าพารามิเตอร์แบบ ML_c (Maximum Likelihood หักเปลี่ยนรูปข้อมูลให้มีการแจกแจงเป็นโค้งปกติ)

2.1 คำนำน้าหนักองค์ประกอบจากการประมาณค่ามีค่าเฉลี่ยต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) และช่วงความเชื่อมั่น 95 % ไม่ครอบคลุมค่าพารามิเตอร์ในทุก ๆ ขนาดกลุ่มตัวอย่าง ส่วนค่าคลาดเคลื่อนมาตรฐานของค่าเฉลี่ยจะมีค่าลดลงเมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้น

2.2 ความแตกต่างระหว่างค่าประมาณพารามิเตอร์และค่าจริงพารามิเตอร์

2.2.1 $n = 80$ พบว่า ในแต่ละคำนำน้าหนักองค์ประกอบมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกคำนำน้าหนักองค์ประกอบ

2.2.2 $n = 160$ พบว่า ในแต่ละคำนำน้าหนักองค์ประกอบมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกคำนำน้าหนักองค์ประกอบ

2.2.3 $n = 370$ พบว่า ในแต่ละคำนำน้าหนักองค์ประกอบมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกคำนำน้าหนักองค์ประกอบ

2.2.4 $n = 623$ พบว่า ในแต่ละคำนำน้าหนักองค์ประกอบมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) ทุกคำนำน้าหนักองค์ประกอบ

2.3 ค่าร้อยละความเอนเอียงสัมพัทธ์ (Relative Bias) ของค่าประมาณพารามิเตอร์

2.3.1 $n = 80$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 6.542 – 15.654 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 3: 6: 0

2.3.2 $n = 160$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 5.252 – 11.241 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 3: 6: 0

2.3.3 $n = 370$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 4.515 – 10.375 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับปานกลาง มีอัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 1: 7: 1

2.3.4 $n = 623$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 4.761 – 9.381 ซึ่งมีค่าอยู่ในระดับปานกลางทุกค่าน้ำหนักองค์ประกอบ ยกเว้น ค่าน้ำหนักองค์ประกอบ A_{B2} ที่มีค่าอยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 8: 1 สรุปในวิธีการ ML ส่วนใหญ่ค่าความเอนเอียงสัมพัทธ์มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นในทุกค่าน้ำหนักองค์ประกอบ ยกเว้น ค่าน้ำหนักองค์ประกอบ A_{A3} , A_{A4} และ A_{B1} ที่ขนาดตัวอย่าง 370 ที่ค่าความเอนเอียงสัมพัทธ์ต่ำกว่าขนาดตัวอย่าง 623

2.4 ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) ของค่าประมาณพารามิเตอร์

2.4.1 $n = 80$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 10.713 – 15.653 มีค่าอยู่ในระดับสูงทุกค่า อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 9: 0: 0

2.4.2 $n = 160$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 7.418 – 10.938 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 3: 6: 0

2.4.3 $n = 370$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 4.614 – 6.922 มีค่าอยู่ในระดับปานกลางทุกค่า ยกเว้น น้ำหนักองค์ประกอบ A_{A4} และ Φ มีค่าอยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 7: 2

2.4.4 $n = 623$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 3.339 – 5.105 มีค่าอยู่ในระดับต่ำทุกค่า ยกเว้น น้ำหนักองค์ประกอบ A_{A3} มีค่าอยู่ในระดับปานกลาง อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 1: 8

สรุป ในวิธีการ ML พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นทุกค่าน้ำหนักองค์ประกอบ และที่ขนาดกลุ่มตัวอย่างเท่ากับ 80 จะให้ค่าที่ไม่เหมาะสม

2.5 ค่าความเชื่อมั่นของตัวแปรแฝง (Construct Reliability: ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (Average Variance Extracted: ρ_v) ของโมเดลการวัด

2.5.1 $n = 80$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

2.5.2 $n = 160$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

2.5.3 $n = 370$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

2.5.4 $n = 623$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

สรุป ในวิธีการ ML พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001 ในทุกขนาดกลุ่มตัวอย่าง แม้ว่าเมื่อกลุ่มตัวอย่างเพิ่มมากขึ้น จะให้ค่าความเชื่อมั่นของโมเดลการวัดที่เข้าใกล้ค่าจริงของพารามิเตอร์ก็ตาม

2.6 ร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit)

เมื่อพิจารณาจากค่าดัชนีวัดความสอดคล้องกลมกลืน (CFI) วิธีการ ML มีอัตราร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ ในขนาดกลุ่มตัวอย่าง 80, 160, 370 และ 623 พบว่า มีค่า 39.0%, 53.0%, 63.0% และ 76.5% เพิ่มขึ้นตามลำดับ และเมื่อพิจารณาจากค่าดัชนี TLI พบว่า มีค่า 49.5%, 62.0%, 76.0% และ 83.0% เพิ่มขึ้นตามลำดับ อย่างไรก็ตาม แม้ว่าร้อยละของโมเดลที่สอดคล้องจะเพิ่มขึ้นตามขนาดของกลุ่มตัวอย่างที่เพิ่มขึ้น แต่ในภาพรวมแล้ว พบว่า ร้อยละของโมเดลที่ไม่สอดคล้องก็ยังคงมีอยู่เป็นจำนวนมากเช่นเดียวกับการใช้วิธีการประมาณค่าพารามิเตอร์แบบ ML ที่ไม่มีการเปลี่ยนรูปข้อมูลให้เป็น โค้งปกติ

3. วิธีการประมาณค่าพารามิเตอร์แบบ WLSMV (Robust Weighted Least Square – Mean and Variance Adjusted χ^2)

3.1 ที่ $n = 80$ พบว่า น้ำหนักองค์ประกอบจากการประมาณค่ามีค่าเฉลี่ยสูงกว่าค่าจริงของพารามิเตอร์ (Overestimate) และช่วงความเชื่อมั่น 95 % บางค่าไม่ครอบคลุมค่าพารามิเตอร์ ที่ $n = 160$ พบว่า น้ำหนักองค์ประกอบจากการประมาณค่ามีค่าเฉลี่ยสูงกว่าค่าจริงของพารามิเตอร์ และเมื่อพิจารณาช่วงความเชื่อมั่น 95 % ส่วนใหญ่ครอบคลุมค่าพารามิเตอร์ มีค่าน้ำหนักองค์ประกอบเพียง 2 ค่า ที่ไม่ครอบคลุมค่าพารามิเตอร์ ส่วนที่ $n = 370$ และ $n = 623$ พบว่า การประมาณค่าพารามิเตอร์มีค่าใกล้เคียงกับค่าจริงของพารามิเตอร์ และช่วงความเชื่อมั่น 95 % ของค่าน้ำหนักองค์ประกอบครอบคลุมค่าพารามิเตอร์ทุกค่า ส่วนค่าคลาดเคลื่อนมาตรฐานของค่าเฉลี่ยจะมีค่าลดลงเมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้น

3.2 ความแตกต่างระหว่างค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์

3.2.1 $n = 80$ พบว่า น้ำหนักองค์ประกอบ A_{B1} มีเพียงค่าเดียว ที่มีความแตกต่างจากค่าจริงของพารามิเตอร์อย่างมีนัยสำคัญทางสถิติ ($p < .001$) โดยที่ค่าประมาณพารามิเตอร์มีค่าสูงกว่าค่าจริงของพารามิเตอร์ (Overestimate) ส่วนค่าน้ำหนักองค์ประกอบอื่น ๆ ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001

3.2.2 $n = 160$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าน้ำหนักองค์ประกอบ

3.2.3 $n = 370$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าน้ำหนักองค์ประกอบ

3.2.4 $n = 623$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าน้ำหนักองค์ประกอบ

3.3 ค่าร้อยละความเอนเอียงสัมพัทธ์ (Relative Bias) ของค่าประมาณพารามิเตอร์

3.3.1 $n = 80$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 5.462 – 14.927 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับปานกลาง ยกเว้น ค่าน้ำหนักองค์ประกอบ A_{B4} และ Φ ที่มีค่าอยู่ในระดับสูง อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 2: 7: 0

3.3.2 $n = 160$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 3.569 – 10.929 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 1: 3: 5

3.3.3 $n = 370$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 2.324 – 6.890 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 2: 7

3.3.4 $n = 623$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 1.776 – 5.626 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับต่ำ อัตราส่วนของระดับสูง: ปานกลาง: ต่ำ = 0: 1: 8

สรุป ในวิธีการ WLSMV ค่าความเอนเอียงสัมพัทธ์ (Relative Bias) มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นในทุกค่าน้ำหนักองค์ประกอบ

3.4 ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) ของค่าประมาณพารามิเตอร์

3.4.1 $n = 80$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 8.089 – 12.530 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง – สูง อัตราส่วนของระดับสูง: ปานกลาง: ต่ำ = 4: 5: 0

3.4.2 $n = 160$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 5.637 – 8.242 มีค่าอยู่ในระดับปานกลางทุกค่า อัตราส่วนของระดับสูง: ปานกลาง: ต่ำ = 0: 9: 0

3.4.3 $n = 370$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 3.606 – 4.988 มีค่าอยู่ในระดับต่ำทุกค่าน้ำหนักองค์ประกอบ อัตราส่วนของระดับสูง: ปานกลาง: ต่ำ = 0: 0: 9

3.4.4 $n = 623$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานมีพิสัยอยู่ในช่วงร้อยละ 2.626 – 3.678 มีค่าอยู่ในระดับต่ำทุกค่าน้ำหนักองค์ประกอบ อัตราส่วนของระดับสูง: ปานกลาง: ต่ำ = 0: 0: 9

สรุป ในวิธีการ WLSMV พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นทุกค่าน้ำหนักองค์ประกอบ

3.5 ค่าความเชื่อมั่นของตัวแปรแฝง (Construct Reliability: ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (Average Variance Extracted: ρ_v) ของโมเดลการวัด

3.5.1 $n = 80$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าสูงกว่าค่าจริงของพารามิเตอร์ (Overestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001

3.5.2 $n = 160$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

3.5.3 $n = 370$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

3.5.4 $n = 623$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

สรุป ในวิธีการ WLSMV พบว่า ที่ $n = 80$ ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าสูงกว่าค่าจริงของพารามิเตอร์ (Overestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001 และที่ขนาดกลุ่มตัวอย่าง 160 ขึ้นไป พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

3.6 ร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit)

เมื่อพิจารณาจากค่าดัชนีวัดความสอดคล้องกลมกลืน (CFI) วิธีการ WLSMV มีอัตราร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ ในขนาดกลุ่มตัวอย่าง 80, 160, 370 และ 623 พบว่า มีค่า 86.0%, 95.0%, 100.0% และ 100.0% เพิ่มขึ้นตามลำดับ และเมื่อพิจารณาจากค่าดัชนี TLI พบว่า มีค่า 90.0%, 95.0%, 100.0% และ 100.0% เพิ่มขึ้นตามลำดับ ซึ่งในภาพรวมแล้ว วิธีการ WLSMV จะให้ผลการทดสอบความสอดคล้องของโมเดลที่ดีที่สุดในทุกขนาดตัวอย่าง

4. วิธีการประมาณค่าพารามิเตอร์แบบ Bayesian

4.1 ที่ขนาดกลุ่มตัวอย่างทุกขนาด พบว่า ช่วงความเชื่อมั่น 95 % ของค่าทำนายองค์ประกอบครอบคลุมค่าจริงของพารามิเตอร์ทุกค่า และมีค่าใกล้เคียงกับค่าจริงของพารามิเตอร์มากขึ้นเมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้น ส่วนค่าคลาดเคลื่อนมาตรฐานของค่าเฉลี่ยจะมีค่าลดลงเมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้น

4.2 ความแตกต่างระหว่างค่าประมาณพารามิเตอร์และค่าจริงพารามิเตอร์

4.2.1 $n = 80$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าทำนายองค์ประกอบ

4.2.2 $n = 160$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าทำนายองค์ประกอบ

4.2.3 $n = 370$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าทำนายองค์ประกอบ

4.2.4 $n = 623$ พบว่า ค่าประมาณพารามิเตอร์และค่าจริงของพารามิเตอร์ ไม่แตกต่างกันที่ระดับนัยสำคัญทางสถิติ .001 ทุกค่าน้ำหนักองค์ประกอบ

4.3 ค่าร้อยละความเอนเอียงสัมพัทธ์ (Relative Bias) ของค่าประมาณพารามิเตอร์

4.3.1 $n = 80$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 5.433 – 17.399 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับปานกลาง ยกเว้น ค่าน้ำหนัก องค์ประกอบ A_{A1} A_{B4} และ Φ ที่มีค่าอยู่ในระดับสูง อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 2: 7: 0

4.3.2 $n = 160$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 3.606 – 12.164 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 1: 3: 5

4.3.3 $n = 370$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 2.505 – 7.798 ส่วนใหญ่มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับต่ำ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 2: 7

4.3.4 $n = 623$ พบว่า ค่าความเอนเอียงสัมพัทธ์มีพิสัยอยู่ในช่วงร้อยละ 1.993 – 6.241 มีค่าความเอนเอียงสัมพัทธ์อยู่ในระดับต่ำทุกค่าน้ำหนักองค์ประกอบ ยกเว้น ค่าน้ำหนักองค์ประกอบ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 1: 8

สรุป ในวิธีการ Bayesian ค่าความเอนเอียงสัมพัทธ์ (Relative Bias) มีค่าลดลง เมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นในทุกค่าน้ำหนักองค์ประกอบ

4.4 ค่าร้อยละความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) ของค่าประมาณพารามิเตอร์

4.4.1 $n = 80$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 8.667 – 12.100 ส่วนใหญ่มีค่าอยู่ในระดับปานกลาง – สูง อัตราส่วน ของระดับ สูง: ปานกลาง: ต่ำ = 5: 4: 0

4.4.2 $n = 160$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 5.895 – 7.983 มีค่าอยู่ในระดับปานกลางทุกค่าน้ำหนักองค์ประกอบ อัตราส่วนของระดับ สูง: ปานกลาง: ต่ำ = 0: 9: 0

4.4.3 $n = 370$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 3.336 – 5.080 มีค่าอยู่ในระดับต่ำทุกค่าน้ำหนักองค์ประกอบ อัตราส่วน ของระดับ สูง: ปานกลาง: ต่ำ = 0: 0: 9

4.4.4 $n = 623$ พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน มีพิสัยอยู่ในช่วงร้อยละ 2.422 – 3.516 มีค่าอยู่ในระดับต่ำทุกค่าน้ำหนักองค์ประกอบ อัตราส่วน ของระดับ สูง: ปานกลาง: ต่ำ = 0: 0: 9

สรุป ในวิธีการ Bayesian พบว่า ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐาน (Relative Standard Error Bias) มีค่าลดลงเมื่อขนาดกลุ่มตัวอย่างเพิ่มขึ้นทุกค่าน้ำหนักองค์ประกอบ และที่ขนาดกลุ่มตัวอย่างเท่ากับ 80 จะให้ค่าที่ไม่เหมาะสม

4.5 ค่าความเชื่อมั่นของตัวแปรแฝง (Construct Reliability: ρ_C) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (Average Variance Extracted: ρ_V) ของโมเดลการวัด

4.5.1 $n = 80$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_C) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_V) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

4.5.2 $n = 160$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_C) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_V) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

4.5.3 $n = 370$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_C) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_V) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

4.5.4 $n = 623$ พบว่า ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_C) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_V) มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

สรุป ในวิธีการ Bayesian พบว่า ในทุกขนาดกลุ่มตัวอย่าง ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_C) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_V) ของโมเดลการวัด มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001

4.6 ร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit)

เมื่อพิจารณาจากค่า Posterior Predictive p - value > .05 วิธีการ Bayesian มีอัตราร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ ในขนาดกลุ่มตัวอย่าง 80, 160, 370 และ 623 พบว่า มีค่า 59.0%, 67.0%, 98.0% และ 100.0% เพิ่มขึ้นตามลำดับ สรุปได้ว่าวิธีการ Bayesian ให้ผลการทดสอบความสอดคล้องของโมเดลที่ดีที่สุดเมื่อใช้กลุ่มตัวอย่างขนาดตั้งแต่ 370 ขึ้นไป

อภิปรายผลการวิจัย

การอภิปรายผลในการเปรียบเทียบประสิทธิภาพของการประมาณค่าพารามิเตอร์ด้วยวิธีการประมาณค่าพารามิเตอร์ที่แตกต่างกัน 4 วิธี ได้แก่ ML, ML_{tr} , WLSMV และ Bayesian และขนาดกลุ่มตัวอย่างที่แตกต่างกัน 4 ขนาด ได้แก่ 80, 160, 370 และ 623 หน่วยตัวอย่าง สำหรับโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันลำดับแรก จำนวน 2 องค์ประกอบ องค์ประกอบละ 4 ตัวแปรสังเกตได้ ภายใต้เงื่อนไขข้อมูลประชากรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกติ สามารถอภิปรายได้ตามประเด็นดังต่อไปนี้

1. การประมาณค่าพารามิเตอร์ด้วยวิธี ML ภายใต้การละเมิดข้อตกลงเบื้องต้นเกี่ยวกับข้อมูลแบบต่อเนื่องและการแจกแจงที่เป็นโค้งปกติ และวิธีการ ML_{tr} (วิธีการ ML หลังจากเปลี่ยนรูปข้อมูลให้มีการแจกแจงเป็นโค้งปกติ) ภายใต้การคำนวณโดยใช้เมตริกสหสัมพันธ์แบบเพียร์สัน (Pearson Product Moment Correlation: PPM) พบว่า แสดงผลลัพธ์ไม่แตกต่างกันในทุก ๆ เงื่อนไข และเมื่อพิจารณาตามขนาดกลุ่มตัวอย่าง พบว่า ในทุกขนาดกลุ่มตัวอย่างแสดงค่าน้ำหนักองค์ประกอบจากการประมาณต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติ ($p < .001$) และช่วงความเชื่อมั่น 95 % ไม่ครอบคลุมค่าจริงของพารามิเตอร์ในทุก ๆ น้ำหนักองค์ประกอบ ค่าความเอนเอียงสัมพัทธ์ของค่าประมาณพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐานอยู่ในระดับสูงในกลุ่มตัวอย่างขนาดเล็ก ($n = 80$) แม้ว่าความเอนเอียงสัมพัทธ์จะลดลงเมื่อกลุ่มตัวอย่างเพิ่มขึ้นก็ยังแสดงค่าไม่เหมาะสม ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าต่ำกว่าค่าจริงของพารามิเตอร์ (Underestimate) อย่างมีนัยสำคัญทางสถิติที่ระดับ .001 ส่วนร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit) วิธีการ ML และ ML_{tr} จะให้ค่าร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์เพิ่มขึ้นตามขนาดของกลุ่มตัวอย่าง แต่จะมีโอกาสผ่านเกณฑ์ความสอดคล้องน้อยกว่าวิธีการ WLSMV และวิธีการ Bayesian

สรุปได้ว่า การประมาณค่าพารามิเตอร์ด้วยวิธี ML สำหรับข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกติ ทั้งวิธีการแบบดั้งเดิม คือ ละเมิดข้อตกลงเบื้องต้น หรือวิธีเปลี่ยนรูป (Transform) ข้อมูลให้มีการแจกแจงเป็นโค้งปกติก่อนประมาณค่านั้น ไม่มีความเหมาะสมในโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน สอดคล้องกับ Kline (2005, pp. 178 – 179) ที่กล่าวว่า วิธีการประมาณค่าแบบ ML ในโมเดลสมการเชิงโครงสร้างสำหรับข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกตินั้นจะให้ค่าความเอนเอียงของค่าประมาณพารามิเตอร์และค่าความคลาดเคลื่อนมาตรฐานสูง ให้ค่าความคลาดเคลื่อนประเภทที่ 1 (Type I Error) และสถิติ Chi – square มีค่าสูงขึ้น แนวโน้มในการปฏิเสธโมเดล

สมมติฐานมีอัตราเพิ่มขึ้นด้วย โดยเฉพาะในกรณีกลุ่มตัวอย่างขนาดเล็ก และระดับการแจกแจงไม่เป็นโค้งปกติมากขึ้น สอดคล้องกับ Finney & DiStefano (2006) ที่กล่าวว่า การละเมิดข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงปกติพหุ (Multivariate Normal Distribution) ของข้อมูลที่มีลักษณะจำแนกกลุ่ม (Categorical Data) การประมาณค่าพารามิเตอร์ด้วยวิธี ML จะให้ผลลัพธ์การประมาณค่าที่ไม่มีความเชื่อมั่น เช่น ค่าดัชนีวัดความสอดคล้องกลมกลืนมีค่าต่ำ มีความเอนเอียงของค่าประมาณพารามิเตอร์และค่าความคลาดเคลื่อนมาตรฐาน ผลลัพธ์จากข้อมูลเชิงประจักษ์จะไม่สอดคล้องกับโมเดลสมมติฐาน โดยขึ้นอยู่กับขนาดกลุ่มตัวอย่าง ระดับการแจกแจงที่ไม่เป็นโค้งปกติ (Degree of Non – normality) และจำนวนกลุ่มที่จำแนก (Number of Categories ด้วย และสอดคล้องกับ Hu & Bentler (1998, p. 427 cited in Finney & DiStefano, 2006, p. 273) ที่กล่าวว่า ถ้าข้อมูลมีการแจกแจงที่ไม่เป็นโค้งปกติในระดับปานกลางถึงมาก และกลุ่มตัวอย่างมีขนาดเล็ก ($n \leq 250$) วิธีการ ML จะให้ค่าดัชนีวัดระดับความกลมกลืน เช่น TLI, CFI หรือ RMSEA มีแนวโน้มที่จะปฏิเสธโมเดลสมมติฐานสูงขึ้น และสอดคล้องกับงานวิจัยของ Trierweiler (2009) ที่ได้ศึกษาประสิทธิภาพการประมาณค่าพารามิเตอร์โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันสำหรับข้อมูลจัดกลุ่มแบบเรียงอันดับในโปรแกรม LISREL โดยการจำลองข้อมูล โดยเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์ 4 แบบ คือ ML, Robust ML, WLS และ Robust DWLS สำหรับโมเดลประชากรที่แตกต่างกัน 2 โมเดล คือ โมเดล 3 องค์ประกอบ 9 ตัวแปรสังเกตได้ และ 3 องค์ประกอบ 18 ตัวแปรสังเกตได้ ระดับการแจกแจงของข้อมูล 5 ระดับ และ ขนาดกลุ่มตัวอย่าง 4 ขนาด ($n = 100, 200, 500$ และ $1,000$) โดยทำการศึกษาเงื่อนไขละ 500 การทดลองซ้ำ (Replications) ผลการศึกษา พบว่า วิธีการ ML (Normal Theory ML) ในเงื่อนไขที่ข้อมูลมีการแจกแจงแบบปกติพหุ แสดงค่าประมาณของพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐานที่มีความเอนเอียงเล็กน้อย ยกเว้นในกลุ่มตัวอย่างขนาดเล็ก ($n = 100$) ที่พบว่า ค่าเบี่ยงเบนมาตรฐาน แสดงค่า Underestimate ในระดับสูง เมื่อระดับการแจกแจงที่ไม่เป็นโค้งปกติมากขึ้นจะแสดงค่าความเอนเอียงสูงขึ้นด้วย โดยเฉพาะในกลุ่มตัวอย่างขนาดเล็ก ขนาดของโมเดลและขนาดกลุ่มตัวอย่างส่งผลต่อความเอนเอียงอย่างมีปฏิสัมพันธ์กัน คือ เมื่อกลุ่มตัวอย่างเพิ่มขึ้นค่าความเอนเอียงจะลดลง และเมื่อจำนวนตัวแปรสังเกตได้ต้ององค์ประกอบมากขึ้นค่าความเอนเอียงของค่าคลาดเคลื่อนมาตรฐานจะลดลง ค่าสถิติ Chi – square และดัชนีวัดความสอดคล้องกลมกลืน คือ CFI และ RMSEA แสดงค่าไม่เหมาะสมในทุก ๆ เงื่อนไข โดยสรุป ในเงื่อนไขข้อมูลจัดกลุ่มแบบเรียงอันดับที่มีการแจกแจงที่ไม่เป็นโค้งปกติ วิธีการ ML ให้ผลการประมาณค่าที่มีประสิทธิภาพต่ำ นอกจากนี้ผลการวิจัยยังสอดคล้องกับการศึกษาอีกหลายงานวิจัย เช่น Bollen (1989), Muthén & Kaplan (1985), Babakus, Ferguson, & Jöreskog (1987), West, Finch, & Curran (1995), Green et al. (1997), Hutchinson & Olmos (1998), Finney & DiStefano (2006)

2. การประมาณค่าพารามิเตอร์ด้วยวิธี WLSMV ภายใต้การคำนวณโดยใช้เมตริกซ์สหสัมพันธ์แบบโพลีคอร์ริก (Polychoric Correlation) ในกลุ่มตัวอย่างขนาดเล็ก ($n = 80$) แสดงน้ำหนักองค์ประกอบจากการประมาณค่าสูงกว่าค่าจริงของพารามิเตอร์ (Overestimate) เล็กน้อย แต่เมื่อกลุ่มตัวอย่างเพิ่มขึ้น ($n = 160, 370$ และ 623) พบว่า น้ำหนักองค์ประกอบจากการประมาณค่าไม่แตกต่างกันจากค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ ($p < .001$) ช่วงความเชื่อมั่น 95 % ครอบคลุมค่าจริงของพารามิเตอร์ทุกค่า ค่าความเอนเอียงสัมพัทธ์ของค่าประมาณพารามิเตอร์และค่าคลาดเคลื่อนมาตรฐานในกลุ่มตัวอย่างขนาดเล็ก ($n = 80$) ส่วนใหญ่อยู่ในระดับปานกลาง แต่เมื่อกลุ่มตัวอย่างเพิ่มขึ้น พบว่า มีความเอนเอียงน้อยลง อยู่ในระดับต่ำ ส่วนค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าสูงกว่าค่าจริงของพารามิเตอร์ (Overestimate) อย่างมีนัยสำคัญทางสถิติ ($p < .001$) ในกลุ่มตัวอย่างขนาดเล็ก แต่เมื่อกลุ่มตัวอย่างเพิ่มขึ้น พบว่า มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ .001 ส่วนร้อยละของโมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit) วิธีการ WLSMV จะให้ผลการทดสอบความสอดคล้องของโมเดลที่ดีที่สุดในทุกขนาดตัวอย่าง

สรุปได้ว่า การประมาณค่าพารามิเตอร์ด้วยวิธี WLSMV โดยการใช้เมตริกซ์สหสัมพันธ์โพลีคอร์ริก (Polychoric Correlation) สำหรับข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกติ ในเงื่อนไขกลุ่มตัวอย่างขนาดเล็ก ($n = 80$) มีความเหมาะสมในระดับปานกลาง และเมื่อขนาดกลุ่มตัวอย่างตั้งแต่ 160 ขึ้นไป พบว่า มีประสิทธิภาพที่เหมาะสมดี ซึ่งสอดคล้องกับ Muthén & Kaplan, (1985); Potthast, (1993); DiStafano, (2002) ที่การศึกษาประสิทธิภาพของวิธีการประมาณค่าพารามิเตอร์ตามทฤษฎีที่พัฒนาขึ้นเพื่อดำเนินการกับข้อมูลที่มีลักษณะจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) ได้แก่ วิธีการ CVM, หรือ Robust WLS (WLSM, WLSMV) ซึ่งพบว่า ค่าประมาณของพารามิเตอร์มีอำนาจการประมาณค่าต่ำเพียงเล็กน้อย เมื่อการแจกแจงข้อมูลไม่เป็นโค้งปกติในระดับมาก ค่าประมาณของพารามิเตอร์มีความแปรปรวนต่อโมเดลขนาดใหญ่ และกลุ่มตัวอย่างมีขนาดเล็ก และสอดคล้องกับงานวิจัยของ Oranje (2003 cited in Trierweiler, 2009, p. 48) ที่ศึกษาเปรียบเทียบวิธีการประมาณค่า แบบ ML, WLS, WLSM และ WLSMV โดยการจำลองข้อมูลประชากรที่มีลักษณะจัดกลุ่มแบบเรียงอันดับ ผลการศึกษาพบว่า วิธีการ WLSMV นั้นมีความแปรปรวนต่อข้อจำกัดเกี่ยวกับกลุ่มตัวอย่างขนาดเล็ก และโมเดลขนาดใหญ่ได้ดีกว่าวิธีการอื่น ๆ และสอดคล้องกับงานวิจัยของ Flora & Curran (2004, p. 466) ที่ศึกษาการประมาณค่าพารามิเตอร์ด้วยเทคนิคการประมาณค่าแบบ WLS และ Robust WLS สำหรับโมเดล CFA ที่ตัวแปรมีมาตรวัดระดับเรียงอันดับ โดยวิธีการจำลองข้อมูล และกำหนด

เงื่อนไขที่แตกต่างกัน คือ กลุ่มตัวอย่าง 4 ขนาด ระดับการแจกแจงข้อมูล 5 ระดับ การกำหนดลักษณะเฉพาะของ โมเดล (Model Specification) 4 แบบ และจำนวน Categories ของตัวแปร 2 แบบ ซึ่งทำการทดลองซ้ำ (Replicate) 500 ครั้ง ผลการศึกษา พบว่า การใช้เมตริกซ์สหสัมพันธ์โพลีคอร์ริก (Polychoric Correlation) มีความแข็งแกร่งต่อการละเมิดข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงแบบปกติ วิธีการ Robust WLS แสดงผลลัพธ์ที่เหมาะสมในทุก ๆ เงื่อนไข และมีความเหมาะสมสำหรับการประมาณค่าในเงื่อนไขที่ตัวแปรมีมาตรวัดแบบเรียงอันดับ ในโมเดลขนาดกลางถึงขนาดใหญ่ และกลุ่มตัวอย่างขนาดกลางถึงขนาดเล็ก ซึ่งสอดคล้องกับงานวิจัยของ Beauducel & Herzberg (2006) ได้ศึกษาเปรียบเทียบวิธีการประมาณค่า 2 แบบ คือ ML และ WLSMV ในโมเดล CFA ที่ข้อมูลประชากรมีการแจกแจงเป็นโค้งปกติ ผลการศึกษา พบว่า วิธีการ WLSMV ให้ค่าน้ำหนักองค์ประกอบที่สูง ค่าประมาณพารามิเตอร์มีความแม่นยำมากกว่า มีอัตราร้อยละความสอดคล้องของโมเดล ในอัตราที่สูงกว่าวิธีการอื่น ๆ และแสดงอัตราที่เหมาะสม ในทุกขนาดกลุ่มตัวอย่าง

3. ประมาณค่าพารามิเตอร์ด้วยวิธี Bayesian ในทุกขนาดกลุ่มตัวอย่างแสดงค่าประมาณพารามิเตอร์ไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ ($p < .001$) ทุกค่าน้ำหนักองค์ประกอบ ช่วงความเชื่อมั่น 95 % ของค่าน้ำหนักองค์ประกอบครอบคลุมค่าจริงของพารามิเตอร์ ทุกค่าค่าความเอนเอียงสัมพัทธ์ของค่าประมาณพารามิเตอร์ที่กลุ่มตัวอย่าง $n = 80$ อยู่ในระดับปานกลาง เมื่อกลุ่มตัวอย่างเพิ่มขึ้นค่าความเอนเอียงจะลดลงอยู่ในระดับต่ำ ความเอนเอียงสัมพัทธ์ของค่าคลาดเคลื่อนมาตรฐานที่กลุ่มตัวอย่าง $n = 80$ และ 160 อยู่ในระดับปานกลาง เมื่อกลุ่มตัวอย่างเพิ่มขึ้นค่าความเอนเอียงจะลดลงอยู่ในระดับต่ำ ค่าความเชื่อมั่นของตัวแปรแฝง (ρ_c) และค่าเฉลี่ยความแปรปรวนที่ถูกสกัดได้ (ρ_v) ของโมเดลการวัด มีค่าไม่แตกต่างกับค่าจริงของพารามิเตอร์ที่ระดับนัยสำคัญทางสถิติ ($p < .001$) ในทุกขนาดกลุ่มตัวอย่าง ส่วนร้อยละของ โมเดลที่สอดคล้องกับข้อมูลเชิงประจักษ์ (Percent of Model Fit) วิธีการ Bayesian ให้ผลการทดสอบความสอดคล้องของโมเดลที่ดีที่สุดเมื่อใช้กลุ่มตัวอย่างขนาดตั้งแต่ 370 ขึ้นไป

สรุปได้ว่า การประมาณค่าพารามิเตอร์ด้วยวิธี Bayesian สำหรับข้อมูลจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็น โค้งปกติ มีความเหมาะสมและมีประสิทธิภาพดีที่สุด แม้ในกลุ่มตัวอย่างขนาดเล็ก ซึ่งสอดคล้องกับงานวิจัยของ Asparouhov & Muthén (2010) ที่ศึกษาวิธีการประมาณค่าแบบ Bayesian ในโมเดล CFA ที่มี 1 องค์ประกอบ ในขนาดกลุ่มตัวอย่างที่แตกต่างกัน 6 ขนาด คือ 50 100 200 500 1,000 และ 5,000 ผลการศึกษา พบว่า ในกลุ่มตัวอย่างขนาดเล็ก ($n = 50$ และ $n = 100$) ไม่พบค่าความเอนเอียง ในเงื่อนไขจำนวนตัวแปรสังเกตได้ ในโมเดล พบว่า จำนวนตัวแปรสังเกตได้ทีมากจะแสดงค่าความเอนเอียงสูงกว่าและมีประสิทธิภาพต่ำกว่าจำนวนตัวแปรสังเกตได้น้อย ๆ แต่เมื่อกลุ่มตัวอย่างมีขนาดเพิ่มขึ้นค่าความเอนเอียงจะลดลง และผลลัพธ์แสดงผลดีขึ้น และ Asparouhov & Muthén (2010, pp. 18 – 22) กล่าวถึงการเปรียบเทียบ

วิธีการประมาณค่าแบบ WLSMV และ Bayesian ใน โมเดล SEM ที่ตัวแปรมีลักษณะจำแนกกลุ่ม (Categorical Variable) และข้อมูลขาดหาย (Missing Data) ผลการศึกษา พบว่า วิธี Bayesian ให้ผลลัพธ์การประมาณค่าที่เข้าใกล้ค่าจริงของพารามิเตอร์มากกว่าวิธี WLSMV และ Asparouhov & Muthén (2010, pp. 45 – 47) กล่าวถึงการเปรียบเทียบประสิทธิภาพวิธีการประมาณค่าแบบ Bayesian และ WLSMV ใน Multiple Indicator Growth Modeling สำหรับตัวแปรจำแนกกลุ่ม โดยจำนวนพารามิเตอร์ที่ทำการประมาณค่า 36 พารามิเตอร์ ผลการศึกษา พบว่า วิธีการประมาณค่าพารามิเตอร์แบบ Bayesian มีความแม่นยำมากกว่าวิธีการ WLSMV และใช้เวลาในการคำนวณน้อยกว่าประมาณ 1.5 เท่า และสอดคล้องกับงานวิจัยของ MacKinnon, et al. (2004 cited in Muthén, 2010, p. 6 – 12) ศึกษาเปรียบเทียบวิธีการประมาณค่าแบบ ML และ Bayesian ใน โมเดล SEM ที่ข้อมูลมีการแจกแจงไม่เป็นโค้งปกติ พบว่า วิธีการ ML ในเงื่อนไขละเมิดข้อตกลงเบื้องต้นในการแจกแจงข้อมูลเป็นโค้งปกติแสดงผลลัพธ์ค่าประมาณของพารามิเตอร์ที่ไม่ถูกต้อง ในขณะที่วิธี Bayesian ทำการวิเคราะห์ข้อมูลโดยการยอมให้การแจกแจงภายหลัง (Posterior Distribution) สามารถมีการแจกแจงไม่เป็นโค้งปกติได้ แสดงผลลัพธ์ค่าประมาณของพารามิเตอร์ที่ดีกว่า

4. การประมาณค่าพารามิเตอร์สำหรับ โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน ลำดับแรก ภายใต้เงื่อนไขข้อมูลประชากรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกติ การใช้เมตริกสหสัมพันธ์โพลีคอรริก (Polychoric Correlation: PC) จะมีความเหมาะสมและให้ผลลัพธ์ที่มีความแม่นยำกว่าการใช้เมตริกสหสัมพันธ์แบบเพียร์สัน (Pearson Product Moment Correlation: PPM) สอดคล้องกับงานวิจัยของ Babakus, Ferguson, & Jöreskog (1987) ที่ศึกษาความอ่อนไหว (Sensitive) ของวิธีการประมาณค่าแบบ Maximum Likelihood (ML) ในเงื่อนไขละเมิดข้อตกลงเบื้องต้นเกี่ยวกับมาตรวัด (Measurement Scale) และลักษณะการแจกแจงของข้อมูล โดยการเปรียบเทียบการใช้เมตริกสหสัมพันธ์ 4 แบบ คือ Product – Moment, Polychoric, Spearman’s Roh และ Kendall’s Tau – b ใน โมเดล CFA องค์ประกอบเดียว ตัวแปรสังเกตได้ 4 ตัวแปร ทำการจำลองข้อมูลตัวแปรจำแนกกลุ่มแบบเรียงอันดับ (Ordered Categorical Data) 5 ระดับ ผลการศึกษา พบว่า การใช้เมตริกสหสัมพันธ์โพลีคอรริก (Polychoric) จะให้ค่าประมาณน้ำหนักองค์ประกอบของที่มีความเหมาะสมมากกว่าเมตริกสหสัมพันธ์แบบอื่น ๆ และสอดคล้องกับงานวิจัยของ Hutchinson & Olmos (1998) ที่ศึกษาเปรียบเทียบวิธีการประมาณค่าพารามิเตอร์ แบบ ML และ WLS โดยทำการจำลองข้อมูลประชากรที่ตัวแปรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับ 5 ระดับ และมีการแจกแจงพหุตัวแปรเป็นโค้งปกติ ในเงื่อนไขที่แตกต่างกัน ดังต่อไปนี้ 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) ระดับการแจกแจงตัวแปรของกลุ่มตัวอย่างที่ไม่เป็นโค้งปกติที่ต่างกัน 4 ระดับ โดยวิธีการ ML ใช้เมตริกซ์สหสัมพันธ์แบบเพียร์สัน (PPM) ส่วนวิธีการ WLS ใช้เมตริกซ์สหสัมพันธ์โพลีคอร์ริก (PC) ผลการศึกษา พบว่า วิธีการ WLS โดยใช้เมตริกซ์สหสัมพันธ์โพลีคอร์ริกให้ผลลัพธ์ความกลมกลืนของโมเดลได้ดีกว่าวิธี ML ในเงื่อนไขการแจกแจงข้อมูลแบบเบ้ ค่าสถิติ Chi – square มีความเอนเอียงในวิธี ML เมื่อข้อมูลมีการแจกแจงแบบเบ้ ในขณะที่วิธี WLS มีความเอนเอียงน้อยมาก และสอดคล้องกับงานวิจัยของ Flora & Curran (2004) ที่ศึกษาผลจากการประมาณค่าพารามิเตอร์ด้วยเทคนิคการประมาณค่าแบบ WLS และ Robust WLS สำหรับโมเดล CFA ที่ตัวแปรมีมาตรวัดระดับเรียงอันดับ 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 ผลการศึกษา พบว่า การใช้เมตริกซ์สหสัมพันธ์โพลีคอร์ริก (Polychoric Correlation) มีความแข็งแกร่งต่อการละเมิดข้อตกลงเบื้องต้นเกี่ยวกับการแจกแจงแบบปกติ และตัวแปรมีมาตรวัดแบบเรียงอันดับนี้ประสิทธิภาพเหมาะสมทั้งในโมเดลขนาดกลางถึงขนาดใหญ่ และกลุ่มตัวอย่างขนาดกลางถึงขนาดเล็ก

5. ขนาดกลุ่มตัวอย่างที่เหมาะสมสำหรับโมเดลการวิเคราะห์องค์ประกอบเชิงยืนยัน ลำดับแรก ภายใต้เงื่อนไขข้อมูลประชากรมีลักษณะจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็นโค้งปกตินั้นขึ้นอยู่กับทางเลือกใช้วิธีการประมาณค่าพารามิเตอร์ที่เหมาะสม ในผลการวิจัยครั้งนี้ ได้แก่ วิธีการ WLSMV ขนาดกลุ่มตัวอย่างที่เหมาะสม คือ $n \geq 160$ หรือ ขนาด 20 เท่าของจำนวนตัวแปรสังเกตได้ในโมเดล สอดคล้องกับข้อเสนอแนะของ Kline (2005, p. 178) และ Costello & Osborne (2005 cited in Schumacker & Lomax, 2010, p. 42) และวิธีการ Bayesian ขนาดกลุ่มตัวอย่างที่เหมาะสม คือ $n \geq 80$ หรือ ขนาด 10 เท่า ของจำนวนตัวแปรสังเกตได้ในโมเดล สอดคล้องกับข้อเสนอแนะของ Bentler & Chou (1987 cited in Schumacker & Lomax, 2010, p. 42) และ Raykov & Marcoulides (2006 cited in Chumney, 2012, p. 29) ส่วนการกำหนดขนาดกลุ่มตัวอย่างโดยใช้สูตรของ Cochran (1977 cited in Bartlett, Kotrlik, & Higgins, 2001, pp. 43 – 50) โดยพิจารณาระดับค่าแอลฟา (Alpha Level) สำหรับข้อมูลแบบจำแนกกลุ่ม (Categorical Data) พบว่า ที่ระดับ $\alpha = .05$ และ $\alpha = .01$ ให้ผลลัพธ์ที่มีความเหมาะสมไม่แตกต่างกัน

ข้อเสนอแนะ

ข้อเสนอแนะในการนำผลการวิจัยไปใช้

จากผลการวิจัยในครั้งนี้ ผู้วิจัยขอเสนอแนะการนำผลไปใช้ ดังนี้

1. ผลการวิจัยพบว่า วิธีการ Maximum Likelihood (ML) ทั้งในกรณีที่ไม่มี การเปลี่ยนรูปข้อมูลและกรณีที่มีการเปลี่ยนรูปข้อมูลให้มีการแจกแจงเป็น โกล์ปกติ นั้น เมื่อใช้ประมาณพารามิเตอร์ภายในโมเดลการวัดที่ข้อมูลมีลักษณะเป็นการจำแนกกลุ่มแบบเรียงอันดับ 5 กลุ่มอันดับ เช่น สเกลการประมาณค่า 5 ระดับของ Likert นั้น จะทำให้ได้ค่าประมาณพารามิเตอร์ที่มีความเอนเอียงและมีขนาดค่าที่ต่ำกว่าค่าพารามิเตอร์จริงของประชากร ดังนั้น ในการทำวิจัยที่ใช้เทคนิคการวิเคราะห์องค์ประกอบเชิงยืนยันแบบลำดับแรกที่มี 2 องค์ประกอบ 8 ตัวแปรสังเกตได้ และมีการออกแบบการวัดโดยใช้สเกลการวัดลักษณะนี้ ผู้วิจัยควรเลือกใช้วิธีการประมาณแบบ Bayesian หรือ WLSMV จะมีความถูกต้องเที่ยงตรงมากกว่าวิธีการ ML และหากผู้วิจัยต้องการควบคุมให้มีค่าความเอนเอียงของค่าประมาณพารามิเตอร์อยู่ในระดับต่ำ ผู้วิจัยควรเลือกใช้กลุ่มตัวอย่างที่มีขนาด 160 หน่วยตัวอย่าง ขึ้นไป

2. ขนาดกลุ่มตัวอย่างที่น้อยที่สุดที่เหมาะสมสำหรับ โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันลำดับแรกที่มี 2 องค์ประกอบ 8 ตัวแปรสังเกตได้ ภายใต้เงื่อนไขข้อมูลประชากรมีลักษณะการจำแนกกลุ่มแบบเรียงอันดับที่มีการแจกแจงไม่เป็น โกล์ปกติ (Non – normal Ordered Categorical Data) นั้น วิธีการประมาณค่าพารามิเตอร์แบบ WLSMV จะเหมาะสมสำหรับกลุ่มตัวอย่างที่มีขนาดมากกว่า 160 หน่วยขึ้นไป (หรือขนาด 20 เท่าของจำนวนตัวแปรสังเกตได้ในโมเดล) ในขณะที่วิธีการ Bayesian มีความเหมาะสมกับกลุ่มตัวอย่างขนาด 80 หน่วยตัวอย่างขึ้นไป (ขนาด 10 เท่าของจำนวนตัวแปรสังเกตได้ในโมเดลการวัด)

3. วิธีการที่นักวิจัยนิยมใช้เพื่อแก้ปัญหาข้อมูลที่มีระดับการวัดจำแนกกลุ่มแบบเรียงอันดับที่มีการเบี่ยงเบนจากการกระจายแบบ โกล์ปกติ (Ordered Categorical Non – normality Data) โดยวิธีการเปลี่ยนรูปข้อมูล (Data Transforming) เพื่อให้แต่ละตัวแปรมีการแจกแจงเป็น โกล์ปกติ (Univariate Normal Distribution) และข้อมูลมีลักษณะต่อเนื่อง (Continuous Data) รวมถึงการเพิ่มขนาดกลุ่มตัวอย่างให้มีขนาดใหญ่ แล้วดำเนินการประมาณค่าพารามิเตอร์ด้วยวิธีการ ML นั้น ไม่ได้ช่วยเพิ่มความเที่ยงตรงให้กับผลการประมาณค่าพารามิเตอร์ของโมเดลแต่อย่างใด ดังนั้น ผู้วิจัยควรเลือกใช้วิธีการประมาณค่าพารามิเตอร์ให้เหมาะสมกับธรรมชาติของระดับการวัดตัวแปร และขนาดของกลุ่มตัวอย่าง จะช่วยให้ผลการวิจัยมีความเที่ยงตรงสูงขึ้น

4. ในกรณีที่ผู้วิจัยละเมิดและละเลยข้อตกลงเบื้องต้นของการวิเคราะห์ในด้านความเป็น โกล์ปกติ และระดับการวัดนั้น แม้ว่าจะใช้กลุ่มตัวอย่างที่มีขนาดมากกว่า 623 หน่วยตัวอย่าง ก็ไม่สามารถชดเชยความถูกต้องของผลการวิเคราะห์ได้ ซึ่งการฝ่าฝืนข้อตกลงเบื้องต้นเหล่านี้ จะส่งผลให้ผลการวิจัยมีความเอนเอียงสูง (High Bias) ดังนั้นวัดความกลมกลืนต่าง ๆ มีประสิทธิภาพต่ำ อันจะนำไปสู่การขาดความเที่ยงตรงภายในของการวิจัย (Low Internal Validity) และส่งผลให้

ขาดความเที่ยงตรงภายนอกของการวิจัยอีกด้วย (Low External Validity) ดังนั้น ผู้วิจัยควรเอาใจใส่ต่อการเลือกใช้วิธีการประมาณค่าพารามิเตอร์ให้เหมาะสมกับสเกลการวัดตัวแปรสังเกตได้และระดับการเบี่ยงเบนจากการกระจายแบบปกติ

ข้อเสนอแนะสำหรับการวิจัยครั้งต่อไป

1. ควรมีการศึกษาประสิทธิภาพของวิธีการประมาณค่าพารามิเตอร์ทั้งสามวิธี (ML, WLSMV, Bayesian) ภายใต้ผลกระทบจากเงื่อนไขอื่น ๆ ดังนี้

1.1 เงื่อนไขด้านประเภทของโมเดล (Model Type) ขนาดของโมเดล (Model Size) และความซับซ้อนของโมเดล (Model Complexity) ที่แตกต่างกัน เช่น โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันที่มีจำนวนองค์ประกอบและตัวแปรสังเกตมากขึ้น โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันอันดับสองและอันดับสาม (Second – order CFA, Third – order CFA) โมเดลการวิเคราะห์องค์ประกอบเชิงยืนยันแบบพหุระดับ (Multilevel CFA) หรือ โมเดลการวิเคราะห์ลักษณะหลากหลาย – วิธีหลาย (Multi – traits Multi – methods Technique)

1.2 เงื่อนไขด้านข้อมูลจากการวัดที่มีระดับของการแจกแจงที่ไม่เป็น โค้งปกติ (Level of Non – normality) ที่แตกต่างกัน เริ่มจากข้อมูลที่มีการเบี่ยงเบนจากการกระจายแบบปกติที่น้อยที่สุดไปจนถึงมากที่สุด

1.3 เงื่อนไขด้านจำนวนกลุ่มที่จำแนกของตัวแปรสังเกตได้ (Degree of Categorization of Observed Variable) ที่แตกต่างกัน เช่น ข้อมูลที่มีสเกลการวัดจำแนกกลุ่มอันดับเป็น 3, 4, 5, 6, 7 กลุ่มอันดับ (Categories) เป็นต้น

2. ในการศึกษาครั้งนี้ ผู้วิจัยใช้ข้อมูลจากการเก็บรวบรวมจากข้อมูลจริงแล้วกำหนดสร้างเป็นประชากรเทียม (Pseudo – population) จึงไม่สามารถควบคุมและกำหนดเงื่อนไขของค่าพารามิเตอร์ทุกค่าให้เป็นไปตามที่ต้องการได้ ดังนั้นควรทำวิจัยเพื่อทดสอบเงื่อนไขของการวิจัยในครั้งนี้ซ้ำอีกครั้งโดยใช้การจำลองข้อมูล (Monte Carlo Simulation) เพื่อตรวจสอบยืนยันความเชื่อมั่นและความเที่ยงตรงของข้อสรุปจากผลการวิจัยอีกครั้งหนึ่ง