

หลักการวิเคราะห์โมเดลสมการโครงสร้างพหุระดับโดยใช้ขนาดตัวอย่าง และวิธีการประมาณค่าที่เหมาะสม

The Principle of Multilevel Structural Equation Modeling Analysis by Using Optimal Sample Size and Estimation Methods

ศิริรัตน์ จำแนกสาร^{1*}

¹ศษ.ด. (การวิจัยและประเมินทางการศึกษา) สำนักทะเบียนและวัดผล มหาวิทยาลัยสุโขทัยธรรมาธิราช libsrj@gmail.com

บทคัดย่อ

การวิเคราะห์โมเดลสมการโครงสร้างพหุระดับเป็นเทคนิคการวิเคราะห์ทางสถิติขั้นสูงที่เกิดจากการบูรณาการแนวคิดการวิเคราะห์โมเดลสมการโครงสร้างกับการวิเคราะห์พหุระดับ เพื่อศึกษาอิทธิพลของตัวแปรทำนายหลายระดับที่มีต่อตัวแปรตาม และสามารถศึกษาปฏิสัมพันธ์ข้ามระดับได้ ทำให้มีความครอบคลุมและลึกซึ้งกว่าสถิติวิเคราะห์แบบประเพณีนิยม และช่วยแก้ไขข้อจำกัดในการวิเคราะห์ของโมเดลสมการโครงสร้างและการวิเคราะห์พหุระดับได้ การเลือกใช้สถิติควรคำนึงถึงข้อตกลงเบื้องต้นทางสถิติ การกำหนดขนาดตัวอย่างที่เหมาะสมโดยพิจารณาขนาดตัวอย่างระดับมหภาคหรือระดับกลุ่มเป็นอันดับแรก โดยขนาดตัวอย่างระดับกลุ่มสูงสุดในการวิเคราะห์ควรมีจำนวนมากกว่า 30 กลุ่มขึ้นไป สำหรับขนาดตัวอย่างระดับจุลภาคหรือระดับบุคคลที่เหมาะสมควรมีขนาดตัวอย่างเท่ากับ 10-20 เท่าของจำนวนพารามิเตอร์ สำหรับการประมาณค่าพารามิเตอร์ด้วยวิธีความเป็นไปได้สูงสุด (ML) และวิธีความเป็นไปได้สูงสุดแบบให้สารสนเทศเต็ม (FIML) ใช้ในกรณีที่ขนาดตัวอย่างแต่ละกลุ่มเท่ากัน และข้อมูลมีการแจกแจงแบบปกติ ส่วนการประมาณค่าพารามิเตอร์ด้วยวิธีกึ่งความเป็นไปได้สูงสุดของ Muthen (MUML) วิธีความเป็นไปได้สูงสุดบางส่วน (PML) และวิธีความเป็นไปได้สูงสุดด้วยค่าความคลาดเคลื่อนมาตรฐานที่แกร่งและไคสแควร์ (MLR) ใช้ในกรณีที่จำนวนหน่วยตัวอย่างในแต่ละกลุ่มไม่เท่ากัน และข้อมูลมีการแจกแจงที่ไม่เป็นโค้งปกติ ทั้งนี้ถ้ากลุ่มตัวอย่างมีขนาดใหญ่เพียงพอ การประมาณค่าพารามิเตอร์ด้วยวิธี ML และวิธี MUML จะให้ค่าที่ใกล้เคียงกัน

คำสำคัญ : โมเดลสมการโครงสร้างพหุระดับ, ความสัมพันธ์เชิงสาเหตุพหุระดับ, ขนาดตัวอย่าง, การประมาณค่า

Abstract

Multilevel structural equation model (MSEM) is an advanced statistical analysis technique that is developed by integrated the analysis of the concept of the Structural equation model (SEM) and Hierarchical linear model (HLM). To study the effect of multilevel predictive variables on the dependent variables and to study cross-level interactions, cause to comprehensive and profound than traditional analysis statistics and can help improve limitations in the analysis HLM and SEM analysis. The sample size that is considered first about appropriate with macro-level or group level. The sample size of the highest level of analysis should be more than 30 groups. For the sample size, the micro-level or individual level should be 10-20 times the number of parameters. Parameter estimation by using methods of the Maximum likelihood (ML) and the Full information maximum likelihood (FIML) suitable for balanced group sizes and data has a normal distribution. For methods of Muthen and Muthen's Quasi-maximum likelihood (MUML), Partial maximum likelihood (PML) and the maximum likelihood with robust standard errors and chi-square (MLR) using for unbalanced group sizes and data have a non-normal distribution. However, if the sample size is large, parameter estimation by using ML and MUML methods have similar results.

Keywords : Multilevel Structural Equation Model, Multilevel SEM, MSEM, Multilevel Causal Model, Sample size, Estimation methods

บทนำ

การวิเคราะห์โมเดลสมการโครงสร้างพหุระดับ (Multilevel Structural Equation Model หรือ MSEM) เป็นเทคนิคการประมาณค่าทางสถิติที่เหมาะสมกับธรรมชาติของข้อมูลที่มีความสัมพันธ์เป็นระดับซ้อนกัน (Nested data) โดยข้อมูลที่อยู่ระดับล่างจะซ้อนอยู่ภายใต้ข้อมูลที่อยู่ระดับบน เช่น ข้อมูลระดับนักเรียนสอดแทรกภายในชั้นเรียน ชั้นเรียนสอดแทรกภายในโรงเรียน โรงเรียนสอดแทรกภายในเขตพื้นที่การศึกษา เป็นต้น และให้ความสำคัญกับความผันแปรของตัวแปรทั้งภายในระดับเดียวกันและต่างระดับกัน ซึ่ง MSEM เป็นเทคนิคการวิเคราะห์ทางสถิติขั้นสูงที่เกิดจากการบูรณาการแนวคิดของการวิเคราะห์โมเดลสมการโครงสร้าง (SEM) และการวิเคราะห์พหุระดับ (HLM) (Kanjana-wasee, 2007, pp.193) ซึ่งแต่ละแนวคิดมีข้อดีดังนี้

- 1) การวิเคราะห์โมเดลสมการโครงสร้าง (SEM) เป็นเทคนิคการวิเคราะห์ทางสถิติที่ใช้ศึกษาโครงสร้างความสัมพันธ์เชิงสาเหตุระหว่างตัวแปรทำนายที่มีต่อตัวแปรตามเพื่อตอบคำถามการวิจัยเกี่ยวกับตัวแปรแฝง (Latent variables) ซึ่งมีข้อดีคือสามารถแยกแยะความคลาดเคลื่อนในการวัดออกจากคะแนนจริงได้ ทำให้ผลการวิเคราะห์ถูกต้องมากขึ้น สามารถตอบคำถามการวิจัยได้ทั้งอิทธิพลทางตรง และอิทธิพลทางอ้อม รวมทั้งการวิเคราะห์อิทธิพลส่งผ่าน (Mediation analysis) (Wiratchai, 1999) แต่ข้อจำกัดของ SEM คือวิเคราะห์ข้อมูลได้เพียงระดับเดียวหากข้อมูลเป็นระดับซ้อนกันจะบังคับให้ตัวแปรที่อยู่ต่างระดับกันนำมาวิเคราะห์ให้เสมือนอยู่ในระดับเดียวกันทั้งหมด ส่งผลให้ละเลยต่อการศึกษาปฏิสัมพันธ์ระหว่างตัวแปรที่อยู่ต่างระดับกันหรือไม่สนใจต่อโครงสร้างตามธรรมชาติของข้อมูลที่เป็นระดับซ้อนกัน (Bryk & Raudenbush, 1992; Kanjana-wasee, 2011)
- 2) การวิเคราะห์พหุระดับ (HLM) เป็นเทคนิคทางสถิติที่ใช้ในการวิเคราะห์อิทธิพลของตัวแปรทำนายหลายระดับที่มีต่อตัวแปรตามโดยตัวแปรทำนายมีโครงสร้างเป็นระดับซ้อนกัน (Hierarchical data) ซึ่งมีจุดเด่นคือสามารถวิเคราะห์ข้อมูลที่มีพหุระดับได้ช่วยแก้ปัญหาเชิงเทคนิคในการวิเคราะห์ข้อมูลแบบประเพณีนิยม โดยวิเคราะห์ข้อมูลเพื่อตอบคำถามการวิจัยเกี่ยวกับปฏิสัมพันธ์ระหว่างตัวแปรที่อยู่ต่างระดับกันหรือปฏิสัมพันธ์ข้ามระดับได้ แต่ไม่ให้ความสำคัญต่อโครงสร้างความสัมพันธ์เชิงสาเหตุระหว่างตัวแปร (Kanjana-wasee, 2007) ดังนั้นการวิเคราะห์โมเดลสมการโครงสร้างพหุระดับจึงเป็นเทคนิคการวิเคราะห์ทางสถิติขั้นสูงแนวใหม่มีประโยชน์ต่อการวิจัยในปัจจุบันที่มีความสอดคล้องกับธรรมชาติของข้อมูลของศาสตร์สาขาวิชาต่าง ๆ มากยิ่งขึ้น

การบูรณาการแนวคิดการวิเคราะห์เชิงสาเหตุและการวิเคราะห์พหุระดับ ประกอบด้วย 2 แนวทาง คือ แนวทางแรกเป็นการขยายขอบเขตของการวิเคราะห์เชิงสาเหตุให้สามารถวิเคราะห์ข้อมูลพหุระดับได้ นักวิจัยที่ศึกษาตามแนวคิดนี้คือ Muthen & Muthen (2004) ผู้พัฒนาโปรแกรม Mplus และแนวทางที่สอง เป็นการขยายขอบเขตของการวิเคราะห์พหุระดับให้สามารถวิเคราะห์การถดถอยพหุคูณแบบมีตัวแปรส่งผ่าน (Mediated Multiple Regression -- MMR) และการวิเคราะห์องค์ประกอบ (Factor analysis -- FA) โดยนักวิจัยที่ศึกษาตามแนวคิดนี้ได้แก่ Goldstein (1995) ; Goldstein et al. (1998) ผู้พัฒนาโปรแกรม MLwiN, Bryk and Raudenbush (1992) และ Raudenbush et al. (2004) ผู้พัฒนาโปรแกรม HLM ซึ่งเป็นเทคนิคการวิเคราะห์ข้อมูลพหุระดับ แต่ไม่สามารถวิเคราะห์โมเดลการวิจัยที่มีลักษณะเป็นโมเดลสมการโครงสร้างที่สร้างขึ้นจากทฤษฎีเพื่อแสดงความสัมพันธ์ระหว่างตัวแปรแฝงกับตัวแปรแฝงด้วยกันได้ รวมทั้งยังไม่สามารถวิเคราะห์ความสัมพันธ์ระหว่างตัวแปรแฝงกับตัวแปรสังเกตได้

การวิเคราะห์โมเดลสมการโครงสร้างพหุระดับ ประกอบด้วยหน่วยการวิเคราะห์ (Unit of analysis) ดังนี้

- 1) หน่วยวิเคราะห์ภายในกลุ่ม (Within level) ซึ่งเป็นหน่วยวิเคราะห์ระดับล่างหรือระดับบุคคล เช่น ระดับจุลภาค (Micro-level unit) ระดับส่วนประกอบ (Elementary unit) และภายในหน่วย (Within units) เป็นต้น
- 2) หน่วยวิเคราะห์ระหว่างกลุ่ม (Between level) ซึ่งเป็นหน่วยวิเคราะห์ในระดับบนหรือระดับกลุ่ม เช่น ระดับมหภาค (Macro-level unit) ระดับกลุ่ม (Clusters unit) และระหว่างหน่วย (Between units) เป็นต้น (Wiratchai, 1999, pp. 292-294)

ดังนั้นสามารถเปรียบเทียบความสามารถการวิเคราะห์ด้วยสถิติ 3 ประเภท ประกอบด้วย โมเดลสมการเชิงโครงสร้าง (SEM) โมเดลการวิเคราะห์พหุระดับ (HLM) และโมเดลสมการโครงสร้างพหุระดับ (MSEM) (Wiratchai, 2009) ได้ดังตารางที่ 1

ตารางที่ 1 เปรียบเทียบการวิเคราะห์โมเดลสมการเชิงโครงสร้าง (SEM) โมเดลการวิเคราะห์พหุระดับ (HLM) และโมเดลการวิเคราะห์เชิงสาเหตุพหุระดับ (MSEM)

ประเด็น	SEM	HLM	MSEM
1. การวิเคราะห์พหุระดับ			
1.1 อิทธิพลสุ่มแปรเปลี่ยนตามความชัน (ความชันเป็นตัวแปรตาม)	✓*	✓	✓
ในการวิเคราะห์การถดถอยพหุคูณ			
1.2 อิทธิพลสุ่มแปรเปลี่ยนตามความชัน (ความชันเป็นตัวแปรตาม)	✗	✗	✓
ในการวิเคราะห์อิทธิพลและการวิเคราะห์องค์ประกอบ			
1.3 รูปแบบย่อยอื่น ๆ	✓	✓	✓
2. การวิเคราะห์องค์ประกอบ (ตัวแปรแฝง)			
2.1 การประมาณค่าน้ำหนักองค์ประกอบ	✓	✓	✓**
2.2 ตัวแปรแฝงเป็นตัวแทนของข้อมูลที่สูญหาย	✓	✓	✓
2.3 การวิเคราะห์องค์ประกอบพหุระดับ	✓	✗	✓
3. เทคนิคการประมาณค่า	ML และวิธีอื่น ๆ	Bayesian + ML	MUML และวิธีอื่น ๆ
3.1 ดัชนีวัดระดับความกลมกลืน (Fit indices)	✓	✗	✓
3.2 ดัชนีดัดแปร (Modification Indexes)	✓	✗	✓
3.3 การผ่อนคลายข้อตกลงของความคลาดเคลื่อน	✓	✗	✓
4. การวิเคราะห์อิทธิพล			
4.1 การประมาณค่าอิทธิพล (DI, IE, TE)	✓	✓	✓
4.2 การทดสอบทฤษฎี (Theory testing)	✓	✗	✓
5. การวิเคราะห์ข้อมูลพหุตัวแปร (Multivariate data)	✓	✓	✓

หมายเหตุ: * หมายถึง pseudo balanced approach

** หมายถึง assume known loading

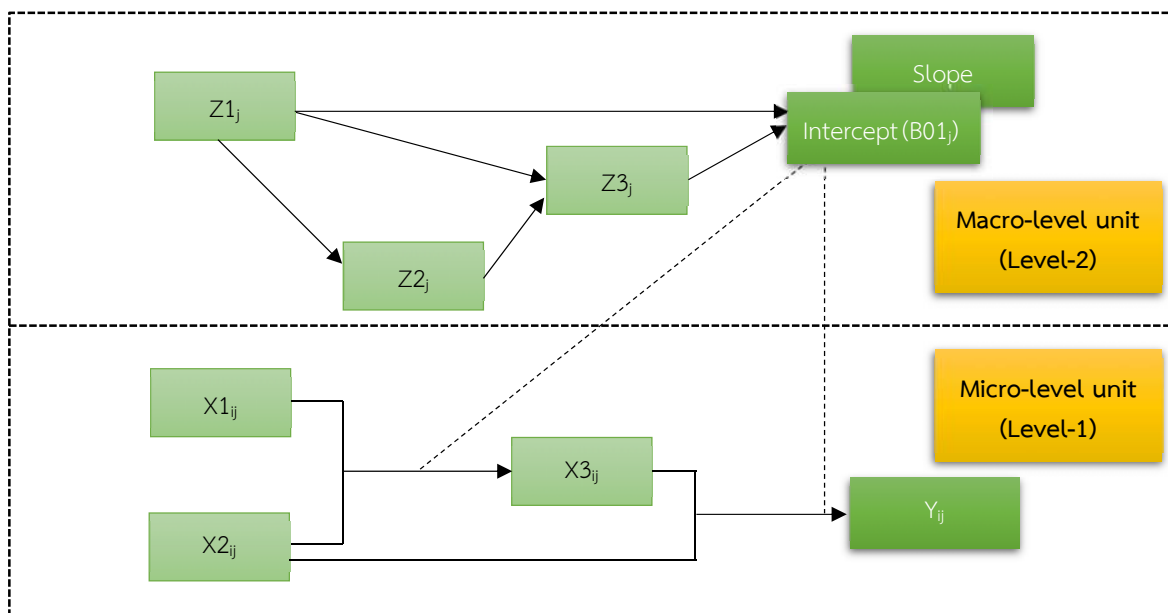
✓ หมายถึง สามารถวิเคราะห์ได้

✗ หมายถึง ไม่สามารถวิเคราะห์ได้

จากการเปรียบเทียบการวิเคราะห์โดยโมเดลสมการเชิงโครงสร้าง (SEM) การวิเคราะห์พหุระดับ (HLM) และโมเดลสมการโครงสร้างพหุระดับ (MSEM) ในตารางข้างต้น จะเห็นได้ว่าโมเดลสมการโครงสร้างพหุระดับ (MSEM) มีจุดเด่นด้านศักยภาพในการวิเคราะห์ข้อมูลที่มีลักษณะเป็นโครงสร้างความสัมพันธ์เชิงสาเหตุและข้อมูลที่มีความสัมพันธ์เป็นระดับซ้อนกันหรือข้อมูลพหุระดับ ซึ่งถือว่ามีศักยภาพในการวิเคราะห์ข้อมูลมากกว่าทั้ง 2 วิธีดังกล่าว ทั้งในเรื่องการวิเคราะห์ที่สามารถสุ่มความชันเข้ามาเป็นตัวแปรตามได้และยังสามารถรวมข้อดีของรูปแบบทั้งสองไว้ในโมเดลสมการโครงสร้างพหุระดับ (MSEM) คือสามารถวิเคราะห์พหุระดับและวิเคราะห์ความสัมพันธ์เชิงสาเหตุได้ (Wiratchai, 2009)

รูปแบบการวิเคราะห์โมเดลสมการโครงสร้างพหุระดับ

เนื่องจากรูปแบบการวิเคราะห์โมเดลสมการโครงสร้างพหุระดับนั้นประกอบด้วยตัวแปรแฝงทั้งในระดับจุลภาค (Micro-level unit) เช่น ระดับบุคคล หรือระดับนักเรียน เป็นต้น และระดับมหภาค (Macro-level unit) เช่น ระดับโรงเรียน เป็นต้น เมื่อนำมาศึกษาพร้อมกันสามารถนำเสนอในรูปแบบภาพโมเดลสมการโครงสร้างพหุระดับพอสังเขปดังภาพต่อไปนี้



ภาพที่ 1 โมเดลสมการโครงสร้างพหุระดับ (2 ระดับ)
ที่มา: Kanjanawasee (2007, pp.148-154)

จากภาพที่ 1 โมเดลสมการโครงสร้างพหุระดับ (2 ระดับ) เป็นตัวอย่างของการศึกษาปัจจัยเชิงสาเหตุที่มีอิทธิพลต่อผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ของนักเรียน ซึ่งจากการศึกษาค้นคว้าทฤษฎี และงานวิจัยที่เกี่ยวข้องทำให้เกิดกรอบแนวคิดการวิจัยว่าปัจจัยสำคัญที่มีอิทธิพลต่อผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์เกิดจากปัจจัยทั้งในระดับนักเรียนและโรงเรียน และได้คัดเลือกตัวแปรเพื่อสร้างโมเดลสมการโครงสร้างพหุระดับดังภาพข้างต้น ซึ่งประกอบด้วยตัวแปรในระดับจุลภาค (Micro-level unit) หรือระดับที่ 1 (Level-1) หรือระดับนักเรียน โดยมีตัวแปรทำนาย 3 ตัวแปรคือ X_{1ij} , X_{2ij} และ X_{3ij} และตัวแปรตาม Y_{ij} หรือผลสัมฤทธิ์ทางการเรียนของนักเรียน i ในโรงเรียน j และตัวแปรระดับมหภาค (Macro-level unit) หรือ Y_j หรือระดับที่ 2 (Level-2) หรือระดับโรงเรียน ประกอบด้วยตัวแปรทำนาย 3 ตัวแปรคือ Z_{1j} , Z_{2j} และ Z_{3j} และตัวแปรตาม (Y_j) หรือค่าเฉลี่ยผลสัมฤทธิ์ทางการเรียนของโรงเรียน j เมื่อสร้างโมเดลสมการโครงสร้างพหุระดับที่เกิดจากทฤษฎีซึ่งแสดงความสัมพันธ์ระหว่างตัวแปรในโมเดลเรียบร้อยแล้ว ซึ่งนำมาเขียนในรูปของสมการที่สำคัญดังนี้

1) การวิเคราะห์ตัวแปรตามด้วยโมเดลไร้มัน (Null Model)

2) การวิเคราะห์โมเดลสมการโครงสร้างระดับต้น โดยวิเคราะห์โมเดลพื้นฐาน (Simple Model) คำนวณค่าสัมประสิทธิ์เส้นทาง ค่า R^2 ทดสอบความสอดคล้องของโมเดลกับข้อมูลเชิงประจักษ์ แปลผลอิทธิพลทางตรงและทางอ้อมของตัวแปรทำนายระดับที่ 1

3) การวิเคราะห์โมเดลสมการโครงสร้างระดับกลาง และระดับสูง โดยวิเคราะห์โมเดลตามสมมุติฐาน (Hypothetical Model) คำนวณค่าสัมประสิทธิ์เส้นทาง ค่า R^2 ทดสอบความสอดคล้องของโมเดลกับข้อมูลเชิงประจักษ์ แปลผลอิทธิพลทางตรงและทางอ้อมของตัวแปรทำนายระดับที่ 2 และระดับที่สูงขึ้นไป (Kanjanawasee, 2007, pp. 148-154)

1. การวิเคราะห์ตัวแปรตามด้วยโมเดลไร้มัน (Null Model)

Level-1: ระดับนักเรียน

$$Y_{ij} = B_{0j} + R_{ij}$$

Level-2: ระดับโรงเรียน

$$B_{0j} = G_{00} + U_j$$

2. วิเคราะห์สมการโครงสร้างระดับที่ 1 สำหรับทำนาย Y_{ij} หรือโมเดลพื้นฐาน (Simple Model)

2.1 วิเคราะห์โมเดลพื้นฐาน สมการ (1)

Level-1: ระดับนักเรียน

$$Y_{ij} = f(X2_{ij}, X3_{ij}, R1_{ij})$$

$$Y_{ij} = B01_j + B21_j(X2_{ij}) + B31_j(X3_{ij}) + R1_{ij}$$

Level-2: ระดับโรงเรียน

$$B01_j = G001 + U01_j$$

$$B21_j = G021 + U21_j$$

$$B31_j = G031 + U31_j$$

2.2 วิเคราะห์โมเดลพื้นฐาน สมการ (2)

Level-1: ระดับนักเรียน

$$X3_{ij} = f(X1_{ij}, X2_{ij}, R2_{ij})$$

$$X3_{ij} = B02_j + B12_j(X1_{ij}) + B22_j(X2_{ij}) + R2_{ij}$$

Level-2: ระดับโรงเรียน

$$B02_j = G002 + U02_j$$

$$B12_j = G012 + U12_j$$

$$B22_j = G022 + U22_j$$

3. วิเคราะห์สมการโครงสร้างระดับที่ 2 สำหรับทำนาย Intercepts $B01_j$ หรือโมเดลตามสมมุติฐาน (Hypothetical

Mode

3.1 วิเคราะห์โมเดลตามสมมุติฐาน สมการ (1)

Level-1: ระดับนักเรียน

$$Y_{ij} = f(X2_{ij}, X3_{ij}, R1_{ij})$$

$$Y_{ij} = B01_j + B21_j(X2_{ij}) + B31_j(X3_{ij}) + R1_{ij}$$

Level-2: ระดับโรงเรียน

$$B01_j = f(Z1_j, Z3_j, U01_j)$$

$$B01_j = G001 + G101(Z1_j) + G301(Z3_j) + U01_j$$

$$B21_j = G021 + U21_j$$

$$B31_j = G031 + U31_j$$

3.2 วิเคราะห์โมเดลตามสมมุติฐาน สมการ (2) โดยใช้โปรแกรมปกติทั่วไป

$$Z3_j = f(Z1_j, Z2_j, U3_j)$$

$$Z3_j = G03 + B13(Z1_j) + G23(Z2_j) + U3_j$$

3.3 วิเคราะห์โมเดลตามสมมุติฐาน สมการ (3) โดยใช้โปรแกรมปกติทั่วไป

$$Z2_j = f(Z1_j, U2_j)$$

$$Z2_j = G02 + G12(Z1_j) + U2_j$$

ข้อตกลงเบื้องต้นทางสถิติ

เนื่องจากสถิติวิเคราะห์แบบประเพณีนิยม เช่น การวิเคราะห์การถดถอยพระระดับ พบว่ามีข้อตกลงเบื้องต้นในการใช้สถิติหลายประการ ดังนี้ 1) ตัวแปรต้นหรือตัวแปรทำนายต้องมีการวัดที่สมบูรณ์ การวัดต้องปราศจากความคลาดเคลื่อน และข้อมูลต้องมีความเป็นอิสระกัน 2) ตัวแปรทำนายแต่ละตัวและตัวแปรตามมีความสัมพันธ์เชิงเส้นตรง 3) ความแปรปรวนของค่าคลาดเคลื่อนมีการแจกแจงแบบปกติ 4) การกระจายของตัวแปรตามในทุกค่าของตัวแปรทำนายมีความแปรปรวนเท่ากัน แต่เนื่องจากธรรมชาติของข้อมูลพระระดับที่มีลักษณะลบล้างกัน ส่งผลทำให้เกิดความไม่เหมาะสมในการเลือกใช้สถิติวิเคราะห์แบบประเพณีนิยมดังกล่าว เพราะจะทำให้ละเมิดข้อตกลงเบื้องต้นโดยเฉพาะประเด็นเกี่ยวกับความเป็นอิสระกัน และการวัดต้องปราศจากความคลาดเคลื่อน สำหรับการวิเคราะห์โมเดลสมการโครงสร้างพระระดับนั้น มีข้อดีในการช่วยผ่อนคลายข้อตกลงเบื้องต้นดังกล่าว ซึ่งยอมให้การวัดมีความคลาดเคลื่อน โดยนำเทอมความคลาดเคลื่อนมาร่วมวิเคราะห์ในโมเดลการวัด และยอมให้ความคลาดเคลื่อนสัมพันธ์กันได้ ซึ่งถือว่าเป็นการผ่อนคลายข้อตกลงเบื้องต้นทางสถิติอีกด้วย

นอกจากนี้การวิเคราะห์โมเดลสมการโครงสร้างพระดับใช้ดัชนีในการตรวจสอบความกลมกลืนของโมเดลกับข้อมูลเชิงประจักษ์ โดยใช้ค่าสถิติไคสแควร์ (Chi-Square Statistics: χ^2) ทั้งนี้ในการทดสอบด้วยสถิติดังกล่าวต้องใช้ด้วยความระมัดระวัง เนื่องจากค่าสถิติ χ^2 มีข้อตกลงเบื้องต้น 4 ประการ คือ 1) ตัวแปรสังเกตได้ภายนอกต้องมีการแจกแจงแบบปกติ 2) การวิเคราะห์ข้อมูลต้องใช้เมทริกซ์ความแปรปรวน-ความแปรปรวนร่วม 3) ขนาดตัวอย่างต้องมีขนาดใหญ่พอ 4) ฟังก์ชันความกลมกลืนมีค่าเป็นศูนย์จริงตามสมมติฐานที่ใช้ในการทดสอบ (Wiratchai, 1999, pp. 53-54)

หลักการพิจารณาตรวจสอบเพื่อการวิเคราะห์โมเดลสมการโครงสร้างพระดับ ประเด็นสำคัญคือการวิเคราะห์สหสัมพันธ์ภายในชั้น (Intraclass correlation -- ICC) เพื่อตรวจสอบว่าตัวแปรต่าง ๆ ในโมเดลการวิเคราะห์มีความผันแปรระหว่างหน่วยเพียงพอกี่จะวิเคราะห์พระดับหรือไม่ โดยถ้าค่า ICC ของทุกตัวแปรควรมีค่ามากกว่าศูนย์จึงถือว่าเหมาะสมที่จะทำการวิเคราะห์พระดับ เกณฑ์การพิจารณาค่า ICC ควรมีค่ามากกว่า 0.05 (Snijders & Bosker, 1999) ถ้าค่า ICC มีค่ามากแสดงว่าตัวแปรมีความสอดคล้องกันสูง แต่ถ้า ICC มีค่าน้อยกว่า 0.05 แสดงว่าข้อมูลในระดับล่างไม่มีความผันแปรในระดับบนจึงไม่จำเป็นที่จะนำข้อมูลไปวิเคราะห์พระดับ

นอกจากนี้อาศัยหลักการพื้นฐานของการวิเคราะห์ SEM ในการพิจารณาดังนี้ 1) ตรวจสอบค่าสหสัมพันธ์ 2) ค่าสถิติ Bartlett's test of sphericity เพื่อทดสอบสมมติฐานว่าเมทริกซ์สหสัมพันธ์นั้นเป็นเมทริกซ์เอกลักษณ์ (Identity matrix) หรือไม่ เนื่องจากโมเดล SEM ประกอบด้วยโมเดลการวัด และโมเดลเชิงโครงสร้าง ซึ่งการวิเคราะห์เพื่อตรวจสอบว่าตัวแปรในโมเดลการวัดมีความเหมาะสมหรือไม่ โดยพิจารณาจากค่าระดับนัยสำคัญทางสถิติที่น้อยกว่าหรือเท่ากับ 0.05 แสดงว่าเมทริกซ์สหสัมพันธ์ของประชากรไม่เป็นเมทริกซ์เอกลักษณ์และเมทริกซ์สหสัมพันธ์นั้นมีความเหมาะสมที่จะใช้วิเคราะห์ต่อไป (Tabachnick & Fidell, 1983; Hair et al., 2010) 3) ค่าดัชนี Kaiser-Meyer-Olkin (KMO) เป็นดัชนีเปรียบเทียบขนาดค่าสัมประสิทธิ์สหสัมพันธ์และขนาดของสหสัมพันธ์บางส่วน (Partial correlation) ระหว่างตัวแปรแต่ละคู่เมื่อขจัดความแปรปรวนของตัวแปรอื่น ๆ ออกไปแล้ว ว่ามีความสัมพันธ์ระหว่างตัวแปรมากพอที่จะนำมาวิเคราะห์องค์ประกอบต่อไปหรือไม่ ถ้าหาก KMO มีค่าเข้าใกล้ 1 แสดงว่ามีความเหมาะสมมาก ส่วนค่าที่น้อยกว่า .50 เป็นค่าที่ไม่เหมาะสมและไม่สามารถยอมรับได้ (Hair et al., 2010)

ขนาดตัวอย่างสำหรับการวิเคราะห์พระดับ

ขนาดตัวอย่างสามารถอธิบายเป็น 2 กลุ่ม คือ ขนาดตัวอย่างระหว่างกลุ่ม (Between Level) และขนาดตัวอย่างภายในกลุ่ม (Within Level) ซึ่งมีรายละเอียดดังนี้

1. ขนาดตัวอย่างระดับมหภาค (Macro-Level) หรือระดับที่ 2 (Level-2) ในการวิเคราะห์โมเดลสมการโครงสร้างพระดับ Muthen (2012) กล่าวถึงขนาดตัวอย่างในระดับกลุ่มควรมีอย่างน้อย 30 – 50 กลุ่ม ส่วน Meuleman and Billiet (2009)

ได้อธิบายว่าจำนวนกลุ่มของการวิเคราะห์พระดับควรมีกลุ่มตัวอย่างอย่างน้อย 40 กลุ่ม เพื่อให้เพียงพอสำหรับการวิเคราะห์ข้อมูลในระดับที่ 2 (Level-2) ซึ่งเป็นการวิเคราะห์อิทธิพลเชิงสาเหตุ (Path analysis) ทั้งนี้หากขนาดอิทธิพลมีเพียงเล็กน้อยจำเป็นต้องใช้กลุ่มตัวอย่างมากกว่า 100 กลุ่ม และทั้งนี้จากการศึกษาพบว่าหากการวิเคราะห์ระดับที่ 2 ซึ่งเป็นโมเดลซับซ้อนแล้วมีกลุ่มตัวอย่างน้อยกว่า 20 กลุ่ม จะส่งผลทำให้เกิดความลำเอียงในการประมาณค่าพารามิเตอร์ได้

นอกจากนี้ ผลงานวิจัยของ Hox and Mass, 2004 (2005); Afshartous and Leeuw (2005); Snijders and Bosker (1999) ได้ข้อค้นพบที่สอดคล้องกันว่าควรให้ความสนใจกับขนาดตัวอย่างที่ใช้ในการวิเคราะห์พระดับโดยเฉพาะที่อยู่ในระดับสูงที่สุด (level-2) มากกว่าระดับที่ต่ำที่สุด (level-1) เพราะจะช่วยลดความคลาดเคลื่อนและเพิ่มความแม่นยำในการประมาณค่าพารามิเตอร์อีกด้วย ผลการวิจัยของ Heck and Thomas (2000); citing Bassiri (1988) อธิบายไว้ว่าขนาดกลุ่มตัวอย่างที่จะทำให้เกิดปฏิสัมพันธ์ข้ามระดับควรมีอย่างน้อย 30 กลุ่ม ผลการวิจัยของ Snijders and Bosker (1999) พบว่ามีข้อเสนอเกี่ยวกับการกำหนดขนาดตัวอย่าง โดยกลุ่มตัวอย่างระดับสูงที่สุดของการวิเคราะห์ควรมีจำนวนมากกว่า 10 กลุ่มขึ้นไป และจากการศึกษาของ Hox and Mass (2004, 2005) มีความเห็นว่าขนาดของตัวอย่างระดับกลุ่มที่สูงที่สุดของการวิเคราะห์ควรจะมีขนาดที่มากกว่าหรือเท่ากับ 30 กลุ่มขึ้นไป

2. ขนาดตัวอย่างระดับจุลภาค (Micro-Level) หรือระดับที่ 1 (Level-1) ในการวิเคราะห์โมเดลสมการโครงสร้างพระดับควรมีกลุ่มตัวอย่างขนาดใหญ่เพียงพอ เนื่องจากใช้ค่าสถิติไคสแควร์ (χ^2) ในการวิเคราะห์ (Wiratchai, 1999, pp. 53-54) โดย Saris and Stronkhorst (1984: 213-214); Anderson and Gerbing (1984: 155-173) ได้ศึกษาพบว่าขนาดตัวอย่างที่เหมาะสมในการวิเคราะห์ในระดับบุคคลของตัวแปรในลักษณะที่มีความสัมพันธ์เชิงสาเหตุควรมีขนาดตัวอย่างไม่ต่ำกว่า 100 หน่วย นอกจากนี้ Lindeman, Merenda and Gold (1980, p. 163) ให้กำหนดอัตราส่วนระหว่างขนาดตัวอย่างต่อจำนวนพารามิเตอร์

หรือตัวแปรว่าควรเป็น 20 หน่วย ต่อ 1 พารามิเตอร์หรือตัวแปรสังเกตได้ ซึ่งสอดคล้องกับ Hair et al. (2010) ได้ให้ข้อพิจารณาขนาดตัวอย่างเท่ากับ 10-20 เท่าของจำนวนพารามิเตอร์

เนื่องจากขนาดตัวอย่างมีความสัมพันธ์กับการประมาณค่าพารามิเตอร์ ถ้าหากกลุ่มตัวอย่างมีขนาดใหญ่การประมาณค่าพารามิเตอร์จะให้ผลการประมาณค่าพารามิเตอร์ที่ใกล้เคียงกัน อย่างไรก็ตามเนื่องจากค่าสถิติไคสแควร์ (χ^2) มีความไวต่อขนาดของกลุ่มตัวอย่าง จึงควรระมัดระวังในการใช้ ค่าสถิติ χ^2 ตัดสินรูปแบบว่ามีความตรงหรือไม่ หรืออีกประการหนึ่งคือ สำหรับกลุ่มตัวอย่างที่มีขนาดใหญ่หรือขนาดตัวอย่างมากกว่า 250 การทดสอบด้วยค่าสถิติ χ^2 จะมีแนวโน้มที่จะปฏิเสธสมมติฐานหลักมากยิ่งขึ้น (Anderson & Gerbing, 1984: 155-173)

การประมาณค่าพารามิเตอร์

การประมาณค่าพารามิเตอร์ในการวิเคราะห์โมเดลสมการโครงสร้างพระระดับ มีหลายวิธีดังนี้

1. วิธีความเป็นไปได้สูงสุด (Maximum likelihood -- ML) หรือวิธีความเป็นไปได้สูงสุดแบบให้สารสนเทศเต็ม (Full information maximum likelihood: FIML) วิธีนี้มีความคงเส้นคงวา มีประสิทธิภาพ และเป็นอิสระจากมาตรวัดใช้ประมาณค่าพารามิเตอร์ในกรณีที่ขนาดตัวอย่างในแต่ละกลุ่มเท่ากัน (Balanced group sizes) และข้อมูลมีการแจกแจงแบบเป็นโค้งปกติ

2. วิธีกึ่งความเป็นไปได้สูงสุดของ Muthen (Muthen and Muthen's Quasi-maximum likelihood -- MUML) หรือวิธีความเป็นไปได้สูงสุดบางส่วน (Partial Maximum Likelihood -- PML) หรือวิธีความเป็นไปได้สูงสุดด้วยค่าความคลาดเคลื่อนมาตรฐานที่แกร่งและไคสแควร์ (Maximum likelihood with robust standard errors and Chi-square: MLR) ใช้ในกรณีจำนวนหน่วยตัวอย่างในแต่ละกลุ่มไม่เท่ากัน (Unbalanced group sizes) และข้อมูลมีการแจกแจงที่ไม่เป็นโค้งปกติ (Wong & Mason, 1985; Goldstein, 1991; Heck & Tomas, 2000; Muthen & Muthen, 2004)

ถ้าหากกลุ่มตัวอย่างมีขนาดใหญ่การประมาณค่าพารามิเตอร์ด้วยวิธี ML และวิธี MUML จะให้ค่าที่ใกล้เคียงกัน สำหรับการแปลงค่าพารามิเตอร์ให้เป็นคะแนนมาตรฐาน (Standardization) โปรแกรม Mplus จะใช้หลักการคำนวณค่าคะแนนมาตรฐานภายในกลุ่มและระหว่างกลุ่ม ซึ่งถ้าหากเป็นการประมาณค่าพารามิเตอร์ของโมเดลภายในกลุ่ม จะพิจารณาที่ค่าความแปรปรวนภายในกลุ่ม และถ้าเป็นการประมาณค่าพารามิเตอร์ของโมเดลระหว่างกลุ่ม จะพิจารณาที่ค่าความแปรปรวนระหว่างกลุ่ม ซึ่งจะเป็นวิธีที่เหมาะสมกับข้อมูลพระระดับ (Muthen, 2012, pp. 259-260)

นอกจากนี้หากจำนวนหน่วยตัวอย่างภายในแต่ละกลุ่มไม่เท่ากัน และตัวแปรมีการแจกแจงไม่เป็นแบบปกติพหุนาม (Multivariate non-normality) จะใช้ฟังก์ชันความกลมกลืน (Fitting function) ในการประมาณค่าพารามิเตอร์ด้วยวิธีความเป็นไปได้สูงสุด (ML) เพื่อให้ค่าความคลาดเคลื่อนมาตรฐานและค่า χ^2 ที่ไม่ลำเอียง (Muthen & Muthen, 2004; Hox, 2002) โดยค่าความคลาดเคลื่อนมาตรฐานของโปรแกรมจะใช้วิธีการประมาณค่าด้วยตัวประมาณเมตริกซ์ความแปรปรวนร่วมที่แกร่ง (Robust covariance matrix estimator) และส่งผลให้ค่าความคลาดเคลื่อนมาตรฐานมีความแกร่ง (Robust standard errors) (Muthen & Muthen, 2004; Freedman, 2006) ส่วนค่า χ^2 สำหรับทดสอบความกลมกลืนประมาณค่าโดยใช้ค่าเฉลี่ยและค่าความแปรปรวนที่ปรับแก้แล้ว (Mean and variance adjustments) ร่วมกับวิธีความเป็นไปได้สูงสุดตามแนวทาง Satorra-Bentler scaled Chi-square (Muthen & Muthen, 2004)

การตรวจสอบความกลมกลืนของโมเดลกับข้อมูลเชิงประจักษ์

การพิจารณาว่าโมเดลที่พัฒนาขึ้นมีความกลมกลืนกับข้อมูลเชิงประจักษ์หรือไม่ โดยใช้วิธีการวิเคราะห์เพื่อตรวจสอบความตรงของรูปแบบซึ่งจะพิจารณาจากค่าสถิติ χ^2 ที่ไม่มีนัยสำคัญทางสถิติ อย่างไรก็ตามการใช้ค่าสถิติ χ^2 ควรใช้ด้วยความระมัดระวังเนื่องจากค่าสถิติ χ^2 มีความไวต่อขนาดของกลุ่มตัวอย่าง และถ้าหากตัวแปรสังเกตได้ข้อมูลมีลักษณะของการกระจายที่ไม่เป็นโค้งปกติ หรือมีจำนวนตัวแปรเชิงกลุ่ม (Categorical data) การทดสอบด้วยค่า χ^2 มีแนวโน้มที่จะปฏิเสธสมมติฐานเช่นกัน (Browne et al., 1984, pp. 62-83) ดังนั้นนักวิจัยจะต้องตัดสินใจด้วยตนเองในการใช้ค่า χ^2 ตรวจสอบความกลมกลืนเพื่อความชัดเจนและถูกต้อง (Bentler & Yuan, 1999, pp.181-197) สำหรับการวัดความกลมกลืนของโมเดลตามกฎพื้นฐาน ให้พิจารณาจากสัดส่วนของค่าสถิติ χ^2 ต่อ df ที่ควรมีค่าน้อยกว่า 2 ($\chi^2/df < 2$) นอกจากนี้ควรพิจารณาความกลมกลืนของโมเดลจากค่าดัชนีต่าง ๆ ดังนี้ (Hox, 2002) (1) ดัชนีวัดระดับความกลมกลืน (GFI) (2) ค่าดัชนีวัดระดับความกลมกลืนที่ปรับแก้แล้ว (AGFI) ซึ่งโดยทั่วไป GFI และ AGFI จะมีค่าระหว่าง 0 ถึง 1 แต่ที่ยอมรับได้ควรมีค่ามากกว่า 0.90

(3) ค่าดัชนีรากของค่าเฉลี่ยกำลังสองของส่วนเหลือ (RMS) (4) ค่าดัชนีรากของค่าเฉลี่ยกำลังสองของส่วนเหลือมาตรฐาน (SRMR) (5) ค่าดัชนีรากของค่าเฉลี่ยกำลังสองของความคลาดเคลื่อน (RMSEA) ซึ่งค่า RMR SRMR และ RMSEA ที่ดีควรมีค่าน้อยกว่า 0.05 จึงจะสรุปได้ว่าโมเดลสอดคล้องกับข้อมูลเชิงประจักษ์ Diamantopoulos & Siguaw (2000) (6) ค่าดัชนี Tucker-Lewis (TLI) ควรมีค่ามากกว่า 0.90 ทั้งนี้สำหรับกลุ่มตัวอย่างที่ไม่เท่ากันควรพิจารณาความกลมกลืนของดัชนี RMSEA และค่า χ^2/df เท่านั้น (Muthen & Muthen, 1998)

ถ้าโมเดลที่ได้ไม่มีความตรงจะปรับโมเดลแล้ววิเคราะห์ใหม่ การปรับแก้โดยใช้ข้อเสนอแนะที่โปรแกรมรายงาน ซึ่งพิจารณาจากดัชนีปรับโมเดล (Modification indices) และพื้นฐานทางทฤษฎีและการวิจัยที่เกี่ยวข้องจนกว่าจะได้โมเดลที่มีความตรง ภายหลังจากที่ได้โมเดลที่มีความตรงแล้วจึงพิจารณาค่าพารามิเตอร์หรือค่าน้ำหนักองค์ประกอบ (Factor loading) ของตัวแปรสังเกตได้ จึงจะทำให้องค์ประกอบที่ต้องการวัดสมบูรณ์และสามารถอธิบายผลได้อย่างถูกต้องแม่นยำ

สรุปผล

การศึกษาโมเดลสมการโครงสร้างพระดับเกิดจากการบูรณาการแนวคิดการวิเคราะห์โมเดลสมการโครงสร้างกับการวิเคราะห์พระดับ โดยมีจุดเด่นด้านความสามารถในการวิเคราะห์ข้อมูลที่มีลักษณะเป็นโครงสร้างความสัมพันธ์เชิงสาเหตุพระดับระหว่างตัวแปรทำนายหลายระดับที่ส่งผลต่อตัวแปรตามโดยสามารถวิเคราะห์ได้โมเดลเดียว นอกจากนี้ในการวิเคราะห์ยังสามารถกลุ่มความชันเข้ามาเป็นตัวแปรตามได้ และยังสามารถวิเคราะห์พระดับและวิเคราะห์ความสัมพันธ์เชิงสาเหตุได้ อย่างไรก็ตามในการเลือกใช้สถิติควรคำนึงถึงข้อตกลงเบื้องต้นทางสถิติ รวมทั้งควรตรวจสอบว่าตัวแปรในโมเดลการวิเคราะห์มีความผันแปรระหว่างหน่วยเพียงพอที่จะวิเคราะห์พระดับหรือไม่ โดยพิจารณาค่าสหสัมพันธ์ภายในชั้น (ICC) ซึ่งควรมีค่ามากกว่า 0.05 รวมทั้งทดสอบเมทริกซ์สหสัมพันธ์นั้นเป็นเมทริกซ์เอกลักษณ์ (Identity matrix) หรือไม่ ด้วยค่าสถิติ Bartlett's test of sphericity และพิจารณาความสัมพันธ์ระหว่างตัวแปรว่ามีมากพอที่จะนำมาวิเคราะห์องค์ประกอบหรือไม่ ด้วยค่าดัชนี Kaiser-Meyer-Olkin (KMO) ถ้ามีค่าเข้าใกล้ 1 แสดงว่ามีความเหมาะสมมาก

ซึ่งการกำหนดขนาดตัวอย่างที่เหมาะสมโดยพิจารณาขนาดตัวอย่างระหว่างกลุ่มเป็นอันดับแรก โดยขนาดตัวอย่างระดับกลุ่มสูงสุดของการวิเคราะห์ควรมีจำนวนมากกว่า 30 กลุ่มขึ้นไป สำหรับขนาดตัวอย่างระดับภายในหน่วยที่เหมาะสมควรมีขนาดตัวอย่างเท่ากับ 10-20 เท่าของจำนวนพารามิเตอร์ สำหรับการประมาณค่าพารามิเตอร์ด้วยวิธีความเป็นไปได้สูงสุด (ML) และวิธีความเป็นไปได้สูงสุดแบบให้สารสนเทศเต็ม (FIML) ใช้ในกรณีที่ขนาดตัวอย่างในแต่ละกลุ่มเท่ากัน และข้อมูลมีการแจกแจงแบบเป็นโค้งปกติ ส่วนการประมาณค่าพารามิเตอร์ด้วยวิธีกำลังความเป็นไปได้สูงสุดของ Muthen (MUML) วิธีความเป็นไปได้สูงสุดบางส่วน (PML) และวิธีความเป็นไปได้สูงสุดด้วยค่าความคลาดเคลื่อนมาตรฐานที่แกร่งและโคสแควร์ (MLR) ใช้ในกรณีที่ขนาดตัวอย่างในแต่ละกลุ่มไม่เท่ากัน และข้อมูลมีการแจกแจงที่ไม่เป็นโค้งปกติ ทั้งนี้ถ้ากลุ่มตัวอย่างมีขนาดใหญ่เพียงพอ

การประมาณค่าพารามิเตอร์ด้วยวิธี ML และวิธี MUML จะให้ค่าที่ใกล้เคียงกัน เมื่อได้ผลวิเคราะห์ทางสถิติในการตรวจสอบว่าโมเดลที่พัฒนาขึ้นมีความกลมกลืนกับข้อมูลเชิงประจักษ์หรือไม่ โดยตรวจสอบค่าสถิติ χ^2 ที่ไม่มีนัยสำคัญทางสถิติ รวมทั้งพิจารณาดัชนีวัดระดับความกลมกลืนต่าง ๆ เช่น GFI, AGFI และ TLI และค่าดัชนีวัดระดับความคลาดเคลื่อนต่าง ๆ เช่น RMR SRMR และ RMSEA ร่วมพิจารณาด้วย

ข้อเสนอแนะ

การวิเคราะห์โมเดลสมการโครงสร้างพระดับถึงแม้จะมีศักยภาพในการวิเคราะห์ข้อมูลเชิงสาเหตุและข้อมูลพระดับได้ โดยอยู่ภายใต้โมเดลเดียวกัน รวมทั้งมีคุณสมบัติช่วยผ่อนคลายนข้อตกลงเบื้องต้นทางสถิติ โดยยอมให้มีความคลาดเคลื่อนในการวัดและความคลาดเคลื่อนสัมพันธ์กันได้ แต่ในการเลือกใช้สถิติดังกล่าวต้องพิจารณาโดยละเอียดตั้งแต่การเลือกตัวแปรเข้ามาศึกษาในโมเดล โดยตัวแปรดังกล่าวต้องมีทฤษฎีและผลงานวิจัยรองรับชัดเจนว่าส่งผลต่อตัวแปรตามที่สนใจศึกษาอย่างแท้จริง การตรวจสอบข้อมูลที่เก็บรวบรวมมาว่ามีคุณสมบัติเพียงพอที่จะวิเคราะห์โมเดลสมการโครงสร้างพระดับหรือไม่ รวมทั้งพิจารณาขนาดตัวอย่างและวิธีการประมาณค่าที่เหมาะสม เนื่องจากสถิติวิเคราะห์ดังกล่าวใช้สถิติทดสอบ χ^2 ซึ่งควรใช้ด้วยความระมัดระวังเพราะสถิติดังกล่าวมีความไวต่อขนาดของกลุ่มตัวอย่าง ข้อมูลมีการกระจายที่ไม่เป็นโค้งปกติ และมีตัวแปรเชิงกลุ่มในโมเดล ทั้งนี้หากเลยข้อพิจารณาข้างต้นจะส่งผลต่อปัญหาในการวิเคราะห์โดยมีแนวโน้มที่ค่าสถิติ χ^2 จะปฏิเสธสมมติฐานหลักมากขึ้น หรือต้องใช้เวลาในการปรับโมเดลหลายรอบจนกว่าโมเดลจะมีความตรงหรือโมเดลที่พัฒนาขึ้นมีความกลมกลืนกับข้อมูลเชิงประจักษ์

บรรณานุกรม

- Afshartous, D. & de Leeuw, J. (2005). Prediction in multilevel models. *Journal of Educational and Behavioral Statistics*, 30(2), 109–139.
- Anderson, J. & Gerbing, D. W. (1984). The effects of sampling errors on convergence, improper solution and goodness of fit indices for maximum likelihood confirmatory factor analysis. *Psychometrika*, 49, 155-173.
- Bentler, P. M. & Yuan, K. H. (1999). Structural equation modelling with small samples: Test statistics. *Multivariate Behavioral Research*, 34(2), 181–197.
- Browne, W., Goldstein, H., Rashbash, J., Plewis, I., Draper, D., & Yang, M. (1998). *User's guide to MLwiN*. London: Institute of Education.
- Bryk, A. S., & Raudenbush, S. W. (1992). *Hierarchical linear models newbury park*. Calif: Sage Publications.
- Diamantopoulos, A. & Siguaw, A. D. (2000). *Introducing LISREL: A guide for the uninitiated*. London : Sage Publications.
- Goldstein, H. (1991). Non-Linear Multilevel Models, with an Application to Discrete Response Data. *Biometrika*, 78, 45-51.
- Goldstein, H. (1995). *Multilevel statistical models*. London : Edward Arnold.
- Hair, J. F., Black, W. C., Babin, B. J., Anderson, R. E., & Tatham, R. L. (1998). *Multivariate data analysis*, 5(3), 207-219.
- Hair, J. F., Black, W.C., Babin, B.J. & Anderson, R. E. (2010). *Multivariate data analysis* (7th ed.). New York: Pearson.
- Heck, R. H. & Thomas, S. L. (2000). An introduction to multilevel modeling techniques. Mahwah, NJ: Lawrence Erlbaum.
- Hox, J. J. & C. J. M. Maas. (2001). The accuracy of multilevel structural equation modeling with pseudobalanced groups and small samples. *Structural Equation Modeling*, 8(2), 157-174.
- Hox, J. J. & C. J. M. Maas. (2004). Robustness issues in multilevel regression analysis. *Statistica Neerlandica*, 18(2), 127–137.
- Hox, J. J. & C. J. M. Maas. (2005). Sufficient sample sizes for multilevel modeling. *Methodology*, 1(3), 86–92.
- Hox, J. J. (2002). *Multilevel analysis: Techniques and applications*. Mahwah, NJ: Erlbaum.
- Kanjanawasee, Sirichai. (2007). *Multi-level analysis* (4th ed.). Bangkok: Chulalongkorn University Printing House. (in Thai)
- Kanjanawasee, Sirichai. (2011). *Multi-level analysis* (5th ed.). Bangkok: Chulalongkorn University Printing House. (in Thai)
- Lindeman, R. H., Merenda, P. F. & Gold, R. Z. (1980). *Introduction to bivariate and multivariate analysis*. Glenview, Illinois: Scott, Foresman and Company.
- Meuleman, B. & Billiet, J. (2009). A Monte Carlo sample size study: how many countries are needed for accurate multilevel SEM?. *Survey Research Methods*, 3(1), 45-58.
- Muthen, B. O. (1994). Multilevel Covariance Structure Analysis. *Sociological Methods and Research*, 22(1), 376-398.
- Muthen, B. O. (2011). *Applications of causally defined direct and indirect effects in mediation analysis using SEM in Mplus*. Retrieved from statmodel.com/download/causalmediation.pdf
- Muthen, L. K. & Muthen, B. O. (1998). *Mplus technical appendices*. Los Angeles, CA: Muthen and Muthen.
- Muthen, L. K. & Muthen, B. O. (2004). *Mplus user's guide* (3rd ed). Los Angeles, CA: Muthen and Muthen.

