

การวิเคราะห์ตัวแบบสมการโครงสร้าง
การวิเคราะห์ Causation Model, Mediation Model, Moderation Model,
Moderated Mediation Model และ Quadratic Effect Model ด้วย Smart PLS4

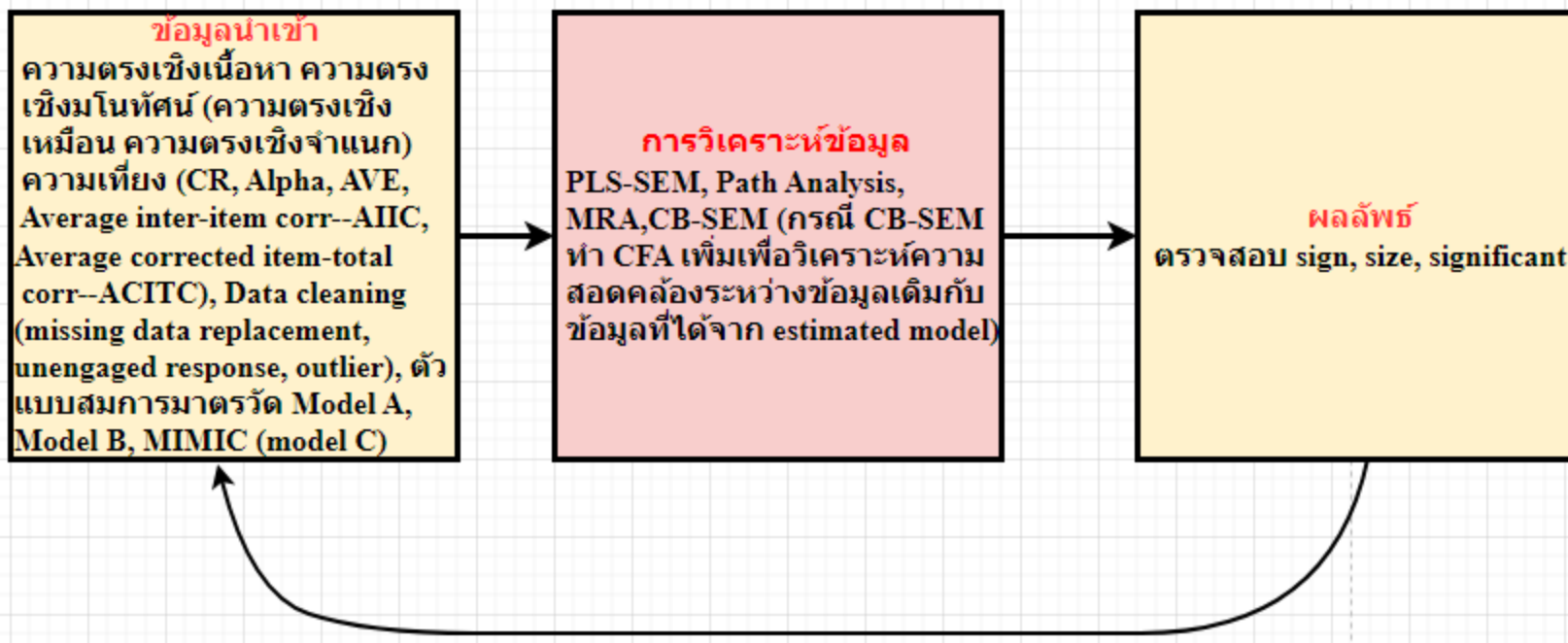
บรรยายในภาคการประชุม Twilight Program ในงาน “มหกรรมวิจัยแห่งชาติ 2568
(Thailand Research Expo 2025)” ณ โรงแรมเซนทาราแกรนด์ และบางกอกคอนเวนชัน
เซ็นเตอร์ เซนทรัลเวิลด์ กรุงเทพมหานคร

วันที่ 19 มิถุนายน 2568

รองศาสตราจารย์ ดร.มนตรี พิริยะกุล

ภาควิชาสถิติ คณะวิทยาศาสตร์ มหาวิทยาลัยรามคำแหง

Data analysis process



Missing data replacement

เรามีวิธีแทนที่ข้อมูลสูญหายหลายวิธี วิธีง่ายคือ

สมมติว่าข้อมูลบันทึกอยู่ใน Excel sheet ซึ่งแถวคือผู้ให้ข้อมูล (case) คอลัมน์คือตัวแปร (item)

1. Mean substitution แทนที่ข้อมูลที่สูญหายด้วยค่าเฉลี่ยข้อมูลของตัวแปรนั้น
2. Mean substitution แทนที่ข้อมูลที่สูญหายด้วยค่าเฉลี่ยข้อมูลของผู้ให้ข้อมูลนั้น
3. Mean of the neighboring cases/items ค่าเฉลี่ยข้อมูลจากผู้ให้ข้อมูลที่อยู่ข้างเคียงซ้ายขวา (แถว) หรือจากข้อมูลของตัวแปรข้างเคียงซ้ายหรือขวาของ construct เดียวกัน (คอลัมน์) จากข้อมูลของผู้ให้ข้อมูล
4. Hot Deck imputation หรือ Similar Case Imputation คือตามหาผู้ให้ข้อมูลคนอื่น ๆ ที่มีคุณลักษณะทางประชากรเหมือนกันกับผู้ที่ไม่ตอบข้อนั้นแล้วใช้ข้อมูลของผู้นั้นแทนที่ข้อมูลที่ไม่ตอบ (คำว่า deck หมายถึง บัตรเจาะรู (punch card) ที่ใช้ในระบบประมวลผลคอมพิวเตอร์สมัยก่อน โดยเราจะตามหาบัตรเจาะรูใบที่เหมือนบัตรที่มีปัญหาข้อมูลสูญหายหรือผิดปกติ คำว่า hot หมายถึงชุด”ปัจจุบัน”
5. Listwise deletion คือทิ้งผู้ให้ข้อมูลไม่ครบ
6. Casewise deletion คือทิ้งข้อมูลเฉพาะรายการ (item) ที่ไม่ตอบของผู้ให้ข้อมูล ข้อมูลรายการอื่นให้รักษาเอาไว้

Unengaged respondent/careless respondent

Unengaged respondent คือผู้ให้ข้อมูลที่ไม่จริง ตอบให้เสร็จ ๆ ไป ลักษณะคำตอบอาจเป็น

1. คำตอบเดิมซ้ำ ๆ เช่น 3,3,3,3,... หรือ 4,4,4,4,... หรือ 5,5,5,5,... เรียกว่า straight-lining
2. คำตอบมีรูปแบบ เช่น 1,2,3,4,5,1,2,3,4,5... เรียกว่า alternative pattern
3. ตอบเสร็จเร็วมากเกินกว่าที่ควรจะเป็น เช่น ควรจะตอบเสร็จประมาณ 5 นาที กลับเสร็จในเวลา 1 -2 นาที
4. คำตอบไม่สอดคล้องกัน เช่น ตัวชี้วัดเรื่องความพึงพอใจในงาน ควรจะล้อกันในระหว่าง sub-domain หรือข้อถามที่คล้าย ๆ กันควรตอบแนวทางเดียวกันกลับตอบขัดกัน เช่น 1. ที่นี่สวยงามมาก 2. ที่นี่ดึงดูดใจมาก กลับตอบขัดกัน
5. Case SD สูงมากหรือต่ำมาก

วิธีตรวจและแก้ปัญหา Unengaged respondent/careless respondent

1. คำนวณหาค่า Case mean, Case SD และ Case CV

SD แสดงการผันแปรของคำตอบของผู้ให้ข้อมูล (case) ถ้า $SD = 0$ แสดงว่าผู้ให้ข้อมูลตอบแบบ straight lining/flatliner ถ้า SD สูงมากแสดงว่าผู้ให้ข้อมูลตอบปะปะแฉ่งทำเป็นว่ามีความถี่ถ้วนในการอ่านคำถามและตอบคำถาม เรียกว่า random response แต่เนื่องจากค่าสูงของ SD คือค่าอนันต์ เราจึงใช้ CV แทน โดยที่ $CV = \frac{SD}{\bar{X}}$ เกณฑ์คือ (Lovie, 2005)

$CV \leq 0.10$ หมายถึง SD ต่ำมาก

$0.10 < CV \leq 0.20$ หมายถึง SD สูงพอดี

$0.20 < CV \leq 0.30$ หมายถึง SD สูงแต่ยอมรับได้

$CV > 0.30$ หมายถึง SD สูงมาก

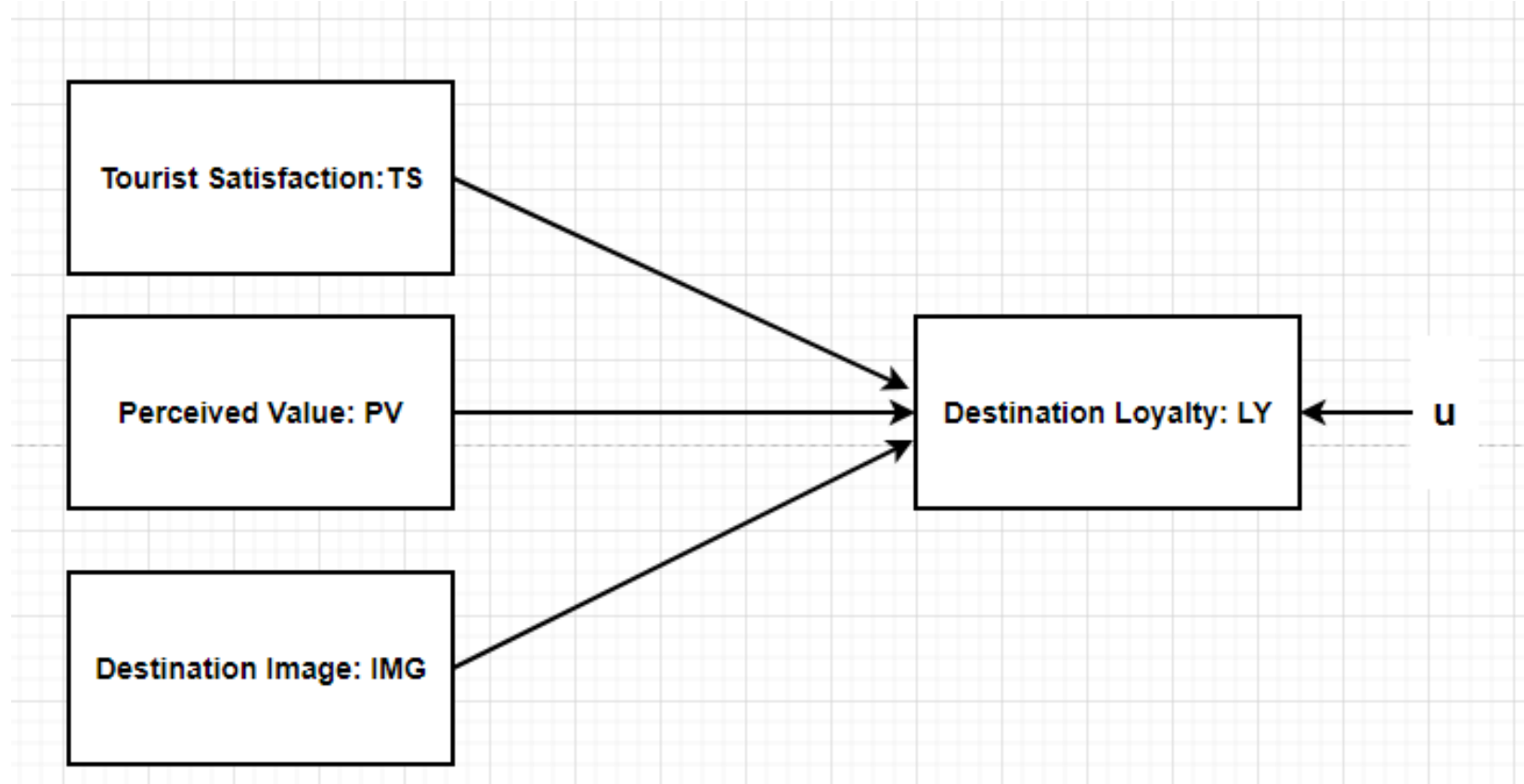
2. ทอยยัด case ที่มี CV ต่ำมากหรือสูงมากทิ้งไป โดยให้ตัดทิ้งเคสหนึ่งรันโปรแกรมครั้งหนึ่งจนกว่าผลลัพธ์จะเหมาะสม คือเครื่องหมายของสัมประสิทธิ์ตรงตามทฤษฎี (sign) มีขนาดค่าของสัมประสิทธิ์ตรงเกณฑ์ (size) และพารามิเตอร์มีนัยสำคัญตรงตามวรรณกรรม (significant)

วิธีตรวจสอบและแก้ปัญหา outlier

Outlier คือข้อมูลที่มีค่าสูงหรือต่ำมาก คืออยู่นอกช่วง 99% confident interval หรือนอก whisker ของ Box plot โดยปกติจะเป็นข้อมูลของ single-valued variable เช่น รายได้ รายจ่าย จำนวนผู้เสียหายจากขบวนการศูนย์รับเรื่อง (call center) อัตราส่วนทางการเงิน เช่น ROA วิธีตรวจสอบที่ง่ายที่สุดก็คือ ถ้าข้อมูลค่าใดมีค่า $|Z| > 3.5$ (คือ $Z > 3.5$ หรือ $Z < -3.5$) โดยที่ Z คือ Z-score นอกจากนี้ยังมี Leverage, Cook's distance, Mahalanobis distance, DFFits (มนตรี พิริยะกุล, 2544) วิธีแก้ปัญหา outlier มีหลายวิธี วิธีที่เรียกว่า Robust method เพราะไม่ถูกระทบจากค่า outlier มีดังนี้

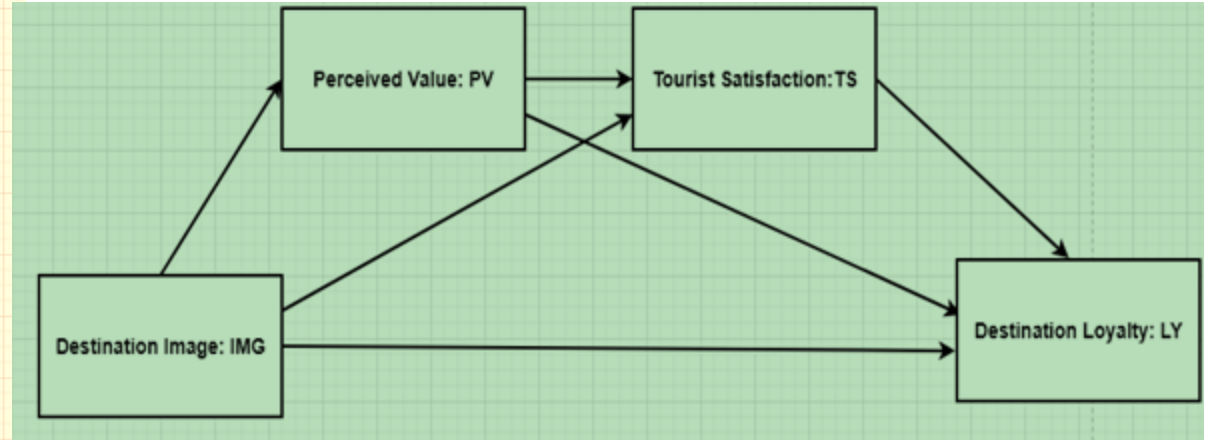
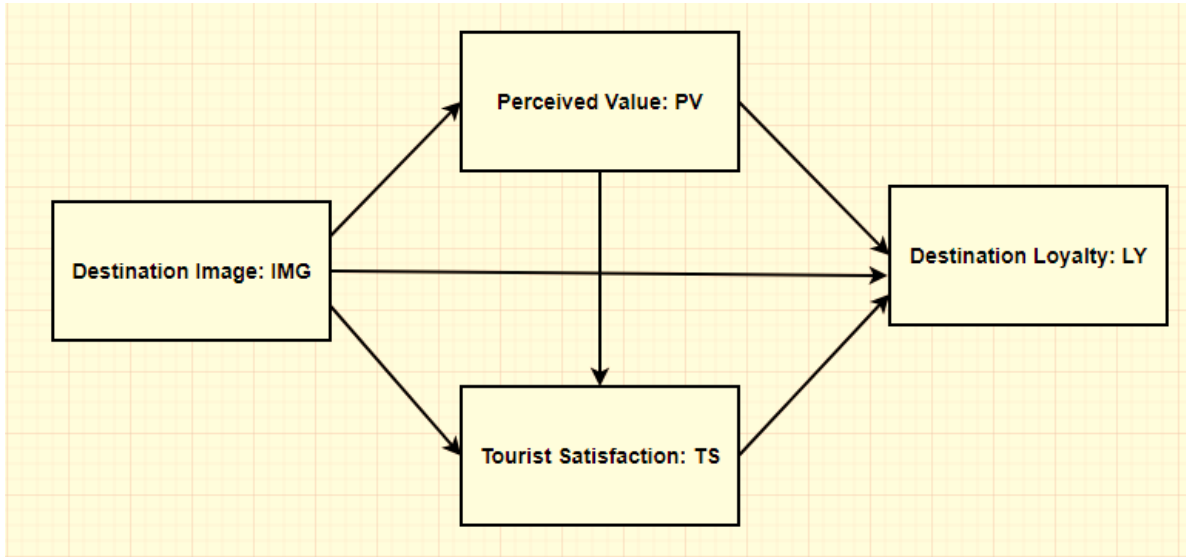
- 1) Winzering แทนที่ค่า outlier ด้วยค่าใกล้ ๆ (nearest non-outlier value)
- 2) ตัดค่า outlier ทิ้ง (Trimming)
- 3) ใช้ median แทน mean ใช้ IQR แทน SD
- 4) IQR normalization ($W = \frac{X - \text{Median}}{\text{IQR}}$)

MRA model



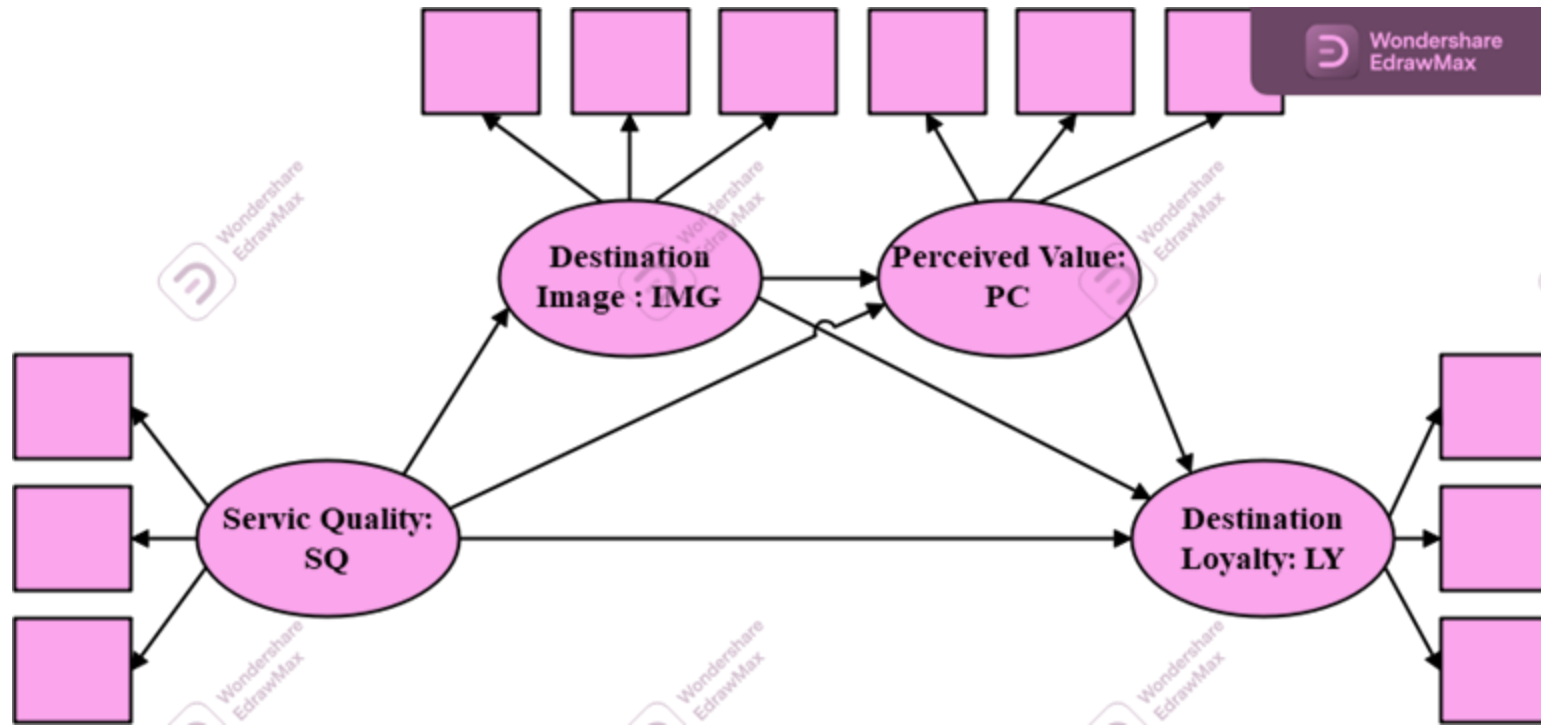
ตัวแบบ Multiple Regression Analysis (MRA) ถือว่าตัวแปรอิสระต้องเป็นอิสระต่อกัน โดยปกติตัวแปรจะเป็น single-valued variable ถ้าตัวแปรอิสระเป็น multi-valued variable /Composite variable จะต้องแปลงค่าให้เป็น single-valued variable โดยทั่วไปจะแปลงเป็นยอดรวม หรือค่าเฉลี่ย หรือ Z-score หรือ factor score

Path model (PA)



MRA model มีข้อตกลงการถดถอยหลายข้อ ข้อหนึ่งคือตัวแปรอิสระต้องเป็นอิสระต่อกัน ถ้าไม่อิสระต่อกันจะเกิดปัญหา multicollinearity ทำให้ผลการวิเคราะห์ผิด สรุปการศึกษาผิด แต่ก็ยอมให้ไม่ต้องเป็นอิสระอย่างสมบูรณ์คือยอมให้ตัวแปรโยงหากันได้ตามทฤษฎีที่เกี่ยวข้อง การกระทำทั้งนี้จึงเกิดเป็น PA model ข้างบนนี้ ภาพขวามือคือภาพเดียวกันกับภาพซ้ายมือแต่จัดภาพให้เป็น serial mediation model

SEM model



ถ้าประสงค์จะใช้ multi-valued variable ตามมาตรวัดเดิม ตัวแบบก็คือ Structural Equation Model (SEM)

ระบบสมการและการวิเคราะห์ตัวแบบ

1. ความเชื่อมโยงระหว่างตัวแปรแฝงใน SEM และความเชื่อมโยงระหว่างตัวแปรใน Path model เรียกว่าโครงสร้างสมการต่าง ๆ ที่เชื่อมโยงตัวแปร/ตัวแปรแฝงเรียกว่าสมการโครงสร้าง (structural equation/model) ทุกสมการล้วนเป็น MRA model จำนวนสมการมีมากเท่ากับจำนวนตัวแปรที่ถูกถูกลูกศรพุ่งเข้าใส่ (คือ internal endogenous variable กับ endogenous variable)
2. เฉพาะใน SEM ความเชื่อมโยงระหว่างตัวแปรแฝงกับตัวชี้วัดเรียกว่าสมการการวัด (measurement equation/model) ทุกสมการเป็น Simple Regression Model (SRA) จำนวนสมการมีมากเท่ากับจำนวนตัวชี้วัด
3. การวิเคราะห์ตัวแบบ SEM จะมี algorithm ไม่เหมือนกันในระหว่าง PLS (variance-based SEM) AMOS หรือ LISREL (covariance-based SEM) และ Path model เช่น PROCESS macro ผลการวิเคราะห์จึงมักจะแตกต่างกัน กรณีที่ผลไม่ตรงกันให้เลือกใช้ซอฟต์แวร์ตัวใดตัวหนึ่งที่ชอบหรือสะดวกใช้ ไม่ควรใช้ซอฟต์แวร์ปนกัน

การออกแบบการวิจัย

1. Causality analysis หรือ Causation มุ่งย้ำความสัมพันธ์ตามเส้นทางที่เป็นจริงตามทฤษฎีและผ่านการยืนยันมาแล้วด้วยผลการวิจัยในอดีตในบริบทต่าง ๆ แต่เราประสงค์จะพิสูจน์ว่าเป็นจริงในบริบทเวลา สถานที่ หรือสังคมของเรา คำถามวิจัยคือ “Does X affect Y?”

2. Mediation analysis มุ่งทำความเข้าใจว่าตัวแปรสาเหตุมีอิทธิพลต่อตัวแปรผลลัพธ์อย่างไร ต้องผ่านตัวแปรอื่นใดหรือไม่ เป็นการขุดค้นหาตัวแปรที่แฝงตัวเชื่อมโยงตัวแปรสาเหตุกับตัวแปรผลลัพธ์ ถ้าตัวแปรคั่นกลางนี้ได้มาจากการปฏิบัติงานจริงมิใช่จากทฤษฎีก็จะเกิดเป็นความรู้ใหม่ที่ได้มาจากการปฏิบัติ แต่ถ้ามาจากทฤษฎีก็จะเป็นการย้ำความจริงนั้นในบริบทของเรา คำถามวิจัยคือ “How does X affect Y?”

3. Moderation analysis 1) มุ่งศึกษาว่าอิทธิพลของปัจจัยสาเหตุที่มีต่อปัจจัยผลลัพธ์ขึ้นอยู่กับอิทธิพลของตัวแปรตัวที่ 3 หรือตัวที่ 4 หรือไม่ หรือปฏิสัมพันธ์ของตัวแปรตัวที่ 3 หรือตัวที่ 4 กับตัวแปรสาเหตุมีผลต่อการเปลี่ยนแปลงอิทธิพลคือขนาดของ beta และทิศทาง (+,-) ที่ปัจจัยสาเหตุมีต่อปัจจัยผลลัพธ์หรือไม่ 2) มุ่งหาเงื่อนไขว่าค่าใดหรือช่วงใดของค่าตัวแปรกำกับเหมาะสม/ยอมรับได้ว่า/นำไปปฏิบัติได้ว่าปัจจัยสาเหตุจะมีผลให้ปัจจัยผลลัพธ์เปลี่ยนแปลงไปอย่างไร รุนแรงขึ้นหรืออ่อนแรงลง หรือเปลี่ยนทิศทางความสัมพันธ์ไปอย่างไรจึงจะเหมาะสม คำถามวิจัยคือ “When does X affect Y?”

การออกแบบการวิจัย

4. Moderated mediation analysis คือตัวแบบที่มุ่งทั้งการค้นหาตัวแปรซ่อนเร้นและค้นหาระดับที่เหมาะสมของตัวแปรที่กำกับความสัมพันธ์ระหว่างตัวแปรตามเส้นทางที่สนใจว่าระดับใด/ค่าใดของตัวแปรกำกับมีผลให้อิทธิพลทางตรงหรืออิทธิพลทางอ้อมมีค่าเหมาะสมที่สุด คำถามวิจัยคือ “How and when does X affect Y?”

5. Quadratic Effect คือตัวแบบที่มุ่งศึกษาว่า

1) ตัวแปรสาเหตุสัมพันธ์กับตัวแปรผลลัพธ์เป็นแบบเส้นตรงหรือไม่ หรือว่าเป็นแบบเส้นโค้ง (curvilinear) แบบ $Y = aX^2 + bX + c$ ถ้าเป็นแบบเส้นโค้งจะมีจุดเปลี่ยนโค้งที่ใด (พิกัดของ X ที่จุดเปลี่ยนโค้งคือ $X = \frac{-b}{2a}$ เกิดจาก $\frac{dY}{dX} = 0$)

2) ถ้าเป็นความสัมพันธ์เชิงเส้นโค้งแสดงว่าการเพิ่มค่าของตัวแปรสาเหตุเพียงเล็กน้อยจะมีผลให้ตัวแปรผลลัพธ์มีค่าเพิ่มขึ้นแบบอัตราเร่ง

คำถามวิจัยคือ “How and when does X affect Y, and is there an optimal point (turning point) in these relationships?”

6. Model with controlled variables คือตัวแบบที่กำหนดให้มีตัวแปรควบคุมเพื่อป้องกัน/แก้ปัญหาคอฟาวน์เดอร์ (confounder)

การออกแบบการวิจัย

ตัวอย่าง mediator ในโลกจริงคือ พ่อสื่อ นายหน้า ตัวแทนขายประกัน ล่าม แพลตฟอร์มตลาดออนไลน์ ระบบธนาคารในธุรกรรมโอนเงิน บรรณธิการวารสาร ครูผู้สอน ดีเจ ทนายความ ตัวแบบอาจเป็นดังนี้และ/หรือมีเส้นทางแสดงอิทธิพลทางตรงจากปัจจัยสาเหตุมายังปัจจัยผลลัพธ์ด้วย

คุณลักษณะของกลุ่มที่ต้องการ → พ่อสื่อ (matchmaker) → หาคู่สำเร็จ

คุณภาพสินค้า/บริการ → นายหน้า → ยอดขายสำเร็จ

ความต้องการประกันของลูกค้า → ตัวแทนขายประกัน → การซื้อกรมธรรม์

ปัญหาการสื่อสาร → ล่าม → ความเข้าใจของผู้ฟัง

ความหลากหลายของสินค้า/บริการ → แพลตฟอร์ม → การตัดสินใจซื้อ

ความต้องการโอนเงิน → ระบบธนาคาร → การโอนเงินสำเร็จ

ต้นฉบับ → บรรณธิการ → บทความที่ตีพิมพ์

เนื้อหาบทเรียน → ครูผู้สอน → ความรู้ความเข้าใจของนักเรียน

ความต้องการฟังเพลง → ดีเจ → ความสุขของผู้ฟัง

ข้อเท็จจริงของคดี → ทนายความ → ผลคดี

การออกแบบการวิจัย

หมายเหตุ ตัวแปรควบคุมคือตัวแปรอิสระอื่นที่ไม่ใช่ตัวแปรอิสระเป้าหมายของการวิจัย แต่ถ้าพิจารณาอย่างรอบคอบแล้วเห็นว่าตัวแปรอิสระที่กำหนดนั้นอธิบายตัวแปรตามได้แต่ก็เห็นว่าอาจมีตัวแปรอื่นอีกที่ทำให้ตัวแปรตามผันแปรไปตามตัวแปรนั้นด้วย เช่น ภาพลักษณ์ปลายทาง → ความตั้งใจมาเที่ยวซ้ำ แต่ระดับรายได้ ระยะทางไกลใกล้ ช่วงอายุ ก็มีส่วนให้ความตั้งใจมาเที่ยวซ้ำ อย่างนี้ เรียกว่า confounding คือไม่รู้แน่ชัดว่าอะไรเป็นสาเหตุของความตั้งใจมาเที่ยวซ้ำ วิธีแก้ไขคือ

1) ให้นำตัวแปรระดับรายได้ ระยะทางไกลใกล้ ช่วงอายุมาร่วมเป็นตัวแปรอิสระ

2) ถ้าการวิจัยมีเป้าหมายเจาะจงที่ภาพลักษณ์ปลายทางไม่ใช่ปัจจัยอื่น ดังนี้ทางออกที่สองคือกำหนดให้ระดับรายได้ ระยะทางไกลใกล้ และช่วงอายุเป็นตัวแปรควบคุม การวิเคราะห์ให้ดำเนินการเช่นเดียวกับข้อ 1) แต่ไม่ต้องสนใจ sign, size และ significant สปส. ของตัวแปรควบคุม ให้สนใจเฉพาะตัวแปรหลักว่ามี การเปลี่ยนแปลงไปหรือไม่ อย่างไร นำเสนอผลเป็น 2 สมการ หรือ 2 conceptual framework

การออกแบบการวิจัย

การวิเคราะห์ตัวแบบที่มีตัวแปรควบคุม

1. การวิเคราะห์ตัวแบบ MRA, PA, SEM ที่มีตัวแปรควบคุมให้แยกวิเคราะห์เป็น 2 ครั้ง เมื่อไม่มีตัวแปรควบคุมครั้งหนึ่งกับเมื่อมีตัวแปรควบคุมอีกครั้งหนึ่ง แล้วอภิปรายผลเฉพาะอิทธิพลของตัวแปรหลัก (กรณี MRA) หรือเส้นทางระหว่างตัวแปรโครงสร้าง (กรณี PA และ SEM) ว่าแตกต่างกันไปอย่างไร มากน้อยเพียงใด โดยเสนอเป็นคอลัมน์เมื่อไม่มีตัวแปรควบคุมและเมื่อมีตัวแปรควบคุมกรณี MRA หรือเสนอเป็น 2 ภาพกรณี SEM ตัวแปรควบคุมให้สรุปผลถึงโดยย่อว่ามีผลต่อการเปลี่ยนแปลงหรือไม่ พร้อมระบุขนาดผลกระทบและนัยสำคัญ

2. กรณีตัวแปรควบคุมเป็นตัวแปรกลุ่ม (categorical variable) ถ้ามีมากกว่า 2 กลุ่ม เช่น การศึกษา (ไม่มีการศึกษา ประถมศึกษา มัธยมศึกษา อุดมศึกษา) ให้แปลงเป็นตัวแปรหุ่น (dummy variable) แต่ละหุ่นเป็น zero-one variable มากเท่าจำนวนกลุ่ม (เช่น k กลุ่ม) แต่ส่งเข้าร่วมวิเคราะห์เพียง k-1 กลุ่มเพื่อป้องกันปัญหาการร่วมเส้นตรงพหุ (multicollinearity) กลุ่มที่ไม่ถูกเลือกจะกลายเป็นกลุ่มอ้างอิงไปโดยปริยาย การพิจารณาเลือกกลุ่มอ้างอิงจะมีผลกระทบต่อการศึกษา จึงต้องพิจารณาให้ดี เช่น การศึกษามี 4 กลุ่มคือ ไม่มีการศึกษา ประถมศึกษา มัธยมศึกษา อุดมศึกษา ให้พิจารณาให้ดีเสียก่อนว่ากลุ่มใดควรเป็นกลุ่มอ้างอิงที่เหมาะสม พิจารณาดีแล้วค่อยกำหนด coding scheme เพื่อให้กลุ่มอ้างอิงมีค่าเท่ากับ 0

การออกแบบการวิจัย

3. ถ้าสนใจให้ตัวแปรกลุ่มเป็นตัวแปรอิสระต้องให้ความสนใจในการแปลผลมากกว่าเมื่อเป็นตัวแปรควบคุมที่เราไม่ค่อยโฟกัสที่การแปลผล/ตีความ แต่ดูว่าผลการวิเคราะห์เปลี่ยนแปลงไปหรือไม่ เพียงใด การแปลผลให้แปลทีละกลุ่มโดยให้เปรียบเทียบว่าตัวแปรหุ่นนั้นกระทบตัวแปรตามมีความแตกต่างกันไปตามกลุ่มเพียงใด เมื่อเปรียบเทียบกับกลุ่มอ้างอิง

4. โปรแกรม PROCESS macro จะมีกล่องโต้ตอบชื่อ covariate ไว้สำหรับให้เราระบุตัวแปรควบคุม แต่โปรแกรมจะทำการควบคุมทุกตัวแปรในรอบที่มีบทบาทเป็นตัวแปรตาม (internal endogenous variable) ถ้าประสงค์จะให้มีการควบคุมเฉพาะตัวแปรตามตัวเดียวตามกรอบ ไม่ใช่ตัวแปรอื่นด้วยให้ใช้วิธี customization หรือเลือกไปใช้วิเคราะห์ด้วย Smart pls4 แทน

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

1. ออกแบบ conceptual framework ให้ยึดตามคำถามการวิจัยและมีวัตถุประสงค์ของการวิจัยเป็นวิธีการหรือแนวทางหรือสิ่งที่ต้องทำ (เรียกว่า road map) เพื่อตอบคำถามวิจัย ปัญหาวิจัยจะประมวลปัจจัยที่พัวพันกับปัญหา และคำถามวิจัยจะถามว่าปัจจัยที่เป็นเป้าหมายถูกเกี่ยวข้อง ถูกกำกับ ถูกควบคุมและรับผลกระทบจากปัจจัยเหล่านั้นหรือไม่ อย่างไร ตัวชี้วัดอาจเป็น reflective model หรือ formative model โดยมีสมมุติฐานการวิจัยเป็นคำตอบคาตอบที่ได้รับมาจากผลการทบทวนวรรณกรรม และจะมีผลการวิจัยเป็นคำตอบที่แท้จริง
2. สร้างเครื่องมือวัดค่าตัวแปรในกรอบแนวความคิดการวิจัย อาจสร้างเองหรือใช้เครื่องมือที่มีผู้อื่นได้พัฒนาเอาไว้แล้ว
3. เก็บรวบรวมข้อมูลจากหน่วยวิเคราะห์ที่กำหนด
4. วิเคราะห์ข้อมูลตามตัวแบบด้วยซอฟต์แวร์ที่เหมาะสม
5. ประเมินคุณภาพเครื่องมือด้วยข้อมูลที่ได้รับ (measurement quality assessment)
6. แปลผล ตีความ อภิปรายผล นัยยะ (implication) ข้อเสนอแนะ (suggestion)

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ ประเมินคุณภาพเครื่องมือด้วยข้อมูลที่ได้รับ (measurement quality assessment) ควรทำให้เสร็จสมบูรณ์ก่อนการสำรวจจริง แต่ปัญหาเรื่องเวลาและค่าใช้จ่ายทำให้ต้องนำข้อมูลจากการสำรวจจริงมาใช้ประเมินความตรงและความเที่ยงซึ่งเป็นการดำเนินการภายหลังซึ่งสวนทางกับสิ่งที่ควรเป็น ดังนั้นเราจึงหวังพึ่ง content validity ซึ่งต้องทำก่อนสำรวจเสมอว่ามาตรวัดจะต้องครอบคลุมเนื้อหา ข้อความกระชับไม่ยาว เนื้อหาไม่ยากแก่การเข้าใจ และไม่มากข้อนเกินไปแต่ก็ต้องครอบคลุมเนื้อหา

สิ่งที่ต้องประเมินคุณภาพเครื่องมือด้วยข้อมูลที่ได้รับ (measurement quality assessment) คือ

1. ความตรงเชิงจำแนก (discriminant validity) แนะนำให้ใช้ HTMT

2. ความตรงเชิงเหมือน (convergent validity) โดยหลักการวัดผลเราต้องสร้างมาตรวัดตัวแปรเดียวกันไว้

หลายวิธีเรียกว่า multi-trait multi-method (MTMM) เช่น วัดด้วยแบบสอบถาม (scale/questionnaire/measures) แบบสังเกตการแสดงออกทางกาย การวัดพฤติกรรม การสัมภาษณ์ ถ้าคะแนนจากมาตรวัดเหล่านี้มีสหสัมพันธ์ต่อกันสูงตั้งแต่ 0.50 เป็นต้นไปแสดงว่ามี convergent validity คือวัดสิ่งที่ตั้งใจจะวัดได้จริง ให้ผลสอดคล้องกับวิธีวัดอื่น จึงเหมาะสมสำหรับนำไปใช้งานจริง

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

แต่ในการปฏิบัติจริงเราวัดที่ loading ว่ามีนัยสำคัญหรือไม่ เป็นปริมาณบวกหรือไม่ มีค่าสูงกว่า 0.707 หรือไม่ ซึ่งไม่ใช่นิยามเดิมของ convergent validity แต่เป็น SEM-based convergent validity คือนักวิจัยในวงการ SEM ยอมรับกันอย่างนี้ มีทั้งในตำรา ทั้งในบทความเผยแพร่และในความเข้าใจของ reviewer คำเรียกที่ถูกต้องคือ uni-dimensionality

3. ความเที่ยง (reliability) คือ สัดส่วนของความผันแปรที่มาจากค่าจริง (true score) เทียบต่อความผันแปรรวม (total variance) หรือพูดแบบเข้าใจง่ายคือ เมื่อเทียบกับความผันแปรทั้งหมดที่เราเห็น มีสัดส่วนเท่าไรที่มาจากสิ่งที่เราต้องการวัดจริงๆ

เราวัดความเที่ยงได้ด้วย Cronbach's alpha, rho-a, rho-c, AVE

ถ้า Cronbach's Alpha = 0.80 หมายความว่า 80% ของความผันแปรในคะแนนที่สังเกตได้มาจากค่าจริงของตัวแปรที่ต้องการวัด ส่วน 20% เป็นความผิดพลาดในการวัด

กรอบแนวความคิดการวิจัย (conceptual framework, CF) vs กรอบการวิจัย (research framework, RF)

หมายเหตุ กรอบแนวความคิดการวิจัย (conceptual framework, CF) โดยปกติจะเป็นกรอบของมโนทัศน์/ความคิด/ตัวแปรแฝง/domain ที่เป็นระดับกว้าง และเชื่อมโยงถึงกันตามทฤษฎี เช่น ภาวะผู้นำ การบริหารทรัพยากรมนุษย์ ความยุติธรรมในองค์การ ผลการปฏิบัติงานของพนักงาน เชื่อมโยงถึงกันเป็นกรอบเรียกว่ากรอบแนวความคิด

หลังจากการทบทวนวรรณกรรมและพิจารณาจุดที่ยังขาดองค์ความรู้ระดับลึก (research gap) และเมื่อได้พิจารณาความเป็นไปได้ของการได้มาของข้อมูลและแนวทางการวิเคราะห์ข้อมูลแล้วให้ปรับ CF อยู่ระดับ sub-domain หรือ sub-sub-domain คือเจาะลึกไปยังองค์ประกอบย่อย/หมวดย่อย (sub-sub-domain) ที่เฉพาะเจาะจงมากขึ้น สามารถหาข้อมูลได้ วิเคราะห์ข้อมูลไม่ซับซ้อนเกินไปและเชื่อมโยงกันด้วยทฤษฎีหรือผลการวิจัยปัจจุบัน เช่น เจาะจงและลดสเกลตัวแปรลงเป็น ภาวะผู้นำตามสถานการณ์ กลยุทธ์การรักษาพนักงาน (Employee Retention Strategy) ความยุติธรรมด้านผลตอบแทน (distributive justice) การทำงานที่มอบหมายเสร็จทันกำหนด (Task Completion on Time) เรียกว่ากรอบการวิจัย (Research Framework, RF)

งานวิจัยจะเริ่มจาก CF แล้วค่อย ๆ ปรับระดับของตัวแปรให้ลึกลงเรียกว่า RF

ช่องว่างการวิจัย (Research Gap)

หมายเหตุ ช่องว่างการวิจัยหรือจุดที่ยังขาดองค์ความรู้ระดับลึก (Research Gap) หมายถึงประเด็นศึกษาหรือปัญหาที่ยังขาดการวิจัยหรือการศึกษาอย่างเพียงพอในสาขาวิชาหนึ่ง ๆ ซึ่งอาจเป็นประเด็นที่ยังไม่เคยถูกศึกษาเลย มีการศึกษาแล้วแต่ยังไม่สมบูรณ์หรือไม่ครอบคลุม ผลการวิจัยยังขัดแย้งไม่ลงรอยกัน ไม่เป็นปัจจุบัน ตัวอย่างเช่น

1. ช่องว่างเชิงทฤษฎี ยังไม่มีทฤษฎีที่อธิบายผลกระทบของ Social Media ต่อพฤติกรรมผู้บริโภค
2. ช่องว่างเชิงวิธีการ งานวิจัยส่วนใหญ่ใช้แบบสอบถามในการเก็บข้อมูล แต่ยังขาดการศึกษาด้วยวิธีการสังเกต หรือการทดลอง การสัมภาษณ์เชิงลึก
3. ช่องว่างเชิงประจักษ์ (หมายถึงสิ่งที่ได้เห็น ได้ยิน ได้รับรส ได้กลิ่น ได้สัมผัสทางกาย ได้รับรู้ทางอารมณ์ในโลกจริง) ยังไม่มีงานวิจัยที่ศึกษาผลกระทบของนโยบาย Work from Home ต่อประสิทธิภาพการทำงานภาครัฐ
4. ช่องว่างเชิงบริบท งานวิจัยส่วนใหญ่ศึกษาในต่างประเทศ ยังขาดการศึกษาในบริบทของประเทศไทย
5. ช่องว่างเชิงปฏิบัติ ยังไม่มีงานวิจัยที่นำเสนอแนวทางในการปรับใช้ AI เพื่อเพิ่มประสิทธิภาพในองค์การ

เชิงประจักษ์ (Empirical)

หมายเหตุ คำว่า "เชิงประจักษ์" (Empirical) หมายถึงสิ่งที่เกี่ยวข้องกับข้อมูลหรือหลักฐานที่ได้จากการสังเกต การทดลอง หรือการเก็บข้อมูลในโลกจริง โดยไม่ขึ้นอยู่กับทฤษฎีหรือความคิดเห็นส่วนตัวเพียงอย่างเดียว ดังนั้น ข้อมูลเชิงประจักษ์ (Empirical Data) คือข้อมูลที่ได้มาจากการสังเกตหรือการทดลองที่สามารถวัดได้ และสามารถนำมาวิเคราะห์เพื่อหาข้อสรุปหรือสร้างความรู้ใหม่ที่เกี่ยวข้องกับข้อมูลหรือหลักฐานที่ได้จากการสังเกต การทดลอง หรือการเก็บข้อมูลในโลกจริงนั้น

ข้อมูลเชิงประจักษ์ไม่ขึ้นอยู่กับทฤษฎีหรือความเชื่อส่วนตัวเพียงอย่างเดียวคือไม่ได้มาจากการคิดเอาเอง หรือการอนุมานจากทฤษฎีที่มีอยู่เดิมเท่านั้น แต่ต้องผ่านการทดสอบหรือตรวจสอบด้วยข้อมูลจากโลกความเป็นจริง เพราะมันเน้นการสังเกต การทดลอง และการพิสูจน์ได้ในโลกจริง ซึ่งช่วยให้การวิจัยมีความน่าเชื่อถือและเป็นกลางมากขึ้น

ทฤษฎีและข้อมูลเชิงประจักษ์มักทำงานร่วมกันเพื่อสร้างความรู้ใหม่และพัฒนาความเข้าใจในปรากฏการณ์ต่าง ๆ ทฤษฎีช่วยให้นักวิจัยรู้ว่าควรสังเกตหรือทดลองอะไร ทฤษฎีช่วยอธิบายความหมาย ทฤษฎีช่วยสร้างสมมติฐาน และข้อมูลเชิงประจักษ์ช่วยทดสอบสมมติฐานนั้น

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

หมายเหตุ ตัวแบบการวัดแบบ Formative Measurement Model (FM) หรือ Composite Measurement Model มีลักษณะดังนี้

1. ตัวแปรแฝงนั้นอาจมีหมวดย่อยหรือองค์ประกอบย่อย (sub-domains) ตามทฤษฎี เป็นองค์ประกอบหรือมิติที่แตกต่างกันแต่รวมกันก่อให้เกิดตัวแปรแฝง ถ้าไม่มี sub-domain ก็ยังยึดหลักการเดียวกันคือตัวชี้วัดก่อตัวเป็นตัวแปรแฝง
2. ตัวชี้วัดแต่ละตัวไม่จำเป็นต้องมีความสัมพันธ์กัน และอาจวัดคนละแง่มุมของตัวแปรแฝง
3. ทิศทางของความสัมพันธ์คือจากตัวชี้วัด → ตัวแปรแฝง
4. ถ้าตัดตัวชี้วัดออกไปหนึ่งตัวจะทำให้ความหมายของตัวแปรแฝงเปลี่ยนไป

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

5. การประเมิน FM ให้ยึดหลักเกณฑ์ของ MRA คือ

1) weight ต้องมีนัยสำคัญ

2) VIF มีค่าต่ำกว่า 5

3) weight ควรมีเครื่องหมายบวก แต่จะเป็นเครื่องหมายใดให้พิจารณาจากทฤษฎี

4) การที่ weight มีเครื่องหมายลบไม่ได้บ่งชี้ว่าตัวชี้วัดนั้นผิดแต่อาจเป็นเพราะมีปัญหาการร่วมเส้นตรง พหุ (multicollinearity) ซึ่งมีวิธีป้องกันคือการทำ mean center กับข้อมูลก่อนนำเข้าสู่กระบวนการวิเคราะห์

5) ถ้า weight ติดลบหรือไม่มีนัยสำคัญอาจคงเอาไว้ถ้า (1) ตัวชี้วัดนั้นมีความสำคัญทางทฤษฎีและ (2) loading มีนัยสำคัญ (ก็ต้องพิจารณาทั้ง weight และ loading)

6) Hair, Hult, Ringle และ Sarstedt (2017) แนะนำให้ทิ้งตัวชี้วัดถ้าพบว่า $VIF > 10$

การตัดสินใจเบื้องต้นว่าตัวชี้วัดเป็นแบบ FM หรือ RM ก่อนกระบวนการวิเคราะห์ทางสถิติให้ดูที่ความสัมพันธ์ระหว่างตัวชี้วัด ถ้าตัวชี้วัดไม่สัมพันธ์กันคือเป็นอิสระต่อกันก็ต้องเป็นแบบ FM ตัวอย่างเช่น การซื้อซ้ำ การบอกต่อ การยอมรับราคา การกล่าวถึงในทางบวก ความเชื่อมั่นในคุณภาพ อิสระต่อกันหรือว่าสัมพันธ์กัน? ถ้าเป็นอิสระต่อกันก็เป็น FM ถ้าสัมพันธ์กันก็เป็น RM

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

หมายเหตุ ตัวแบบการวัดแบบ Reflective Measurement Model (RM) เป็นการมองว่าผลหรือการแสดงออก (ตัวชี้วัด) เกิดมาจากตัวแปรแฝง (ตัวแปรแฝง → ผลที่สังเกตได้) มีลักษณะดังนี้

1. ตัวแปรแฝงนั้นเป็นสาเหตุ (cause) ที่ส่งผลให้เกิดตัวชี้วัดหรือสิ่งที่สังเกตได้
2. ตัวชี้วัดทั้งหมดควรมีความสัมพันธ์กันในทางบวกเพราะเกิดจากสาเหตุเดียวกัน
3. ทิศทางของความสัมพันธ์คือ ตัวแปรแฝง → ตัวชี้วัด
4. สามารถตัดตัวชี้วัดบางตัวออกไปได้โดยไม่เปลี่ยนความหมายของตัวแปรแฝง
5. อาจต้องกำหนด sub-domains ขึ้นมาใหม่เพื่อใช้เป็นกรอบหรือแนวทางร่างแบบสอบถาม โดยพิจารณาว่าอะไรคือผลสะท้อนของตัวแปรแฝงนั้น ซึ่งอาจเป็นแต่ละ sub-domain ตามทฤษฎีเกี่ยวกับตัวแปรแฝงนั้น โดยทั่วไป sub-domain ตามทฤษฎีของเรื่องนั้นส่วนมากจะเป็น formative model แต่ก็อาจเป็น reflective model จึงต้องใช้ความรอบรู้ ความรอบคอบ และมุมมองเชิงพฤติกรรม เชิงอารมณ์ความรู้สึกในการพิจารณาด้วย

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

คือเวิร์ดเบื่องต้นอีกประการหนึ่งคือ กรณี RM ให้พิจารณาว่า

“LV ก่อให้เกิด ...” หรือ

“ LV เป็นสาเหตุของ ...” หรือ

“LV ประกอบด้วย ...”

ส่วนกรณี FM ให้พิจารณาว่า

“...ร่วมกันก่อให้เกิด LV” ,

“LV ประกอบขึ้นมาจาก...”

อย่าพิจารณาโดยใช้คำว่า “เพราะ ... จึงทำให้เกิด...” เช่น 1) เพราะมีความรักดีต่อผลิตภัณฑ์มากจึงซื้อซ้ำ 2) เพราะพึงพอใจในงานมากจึงพอใจค่าตอบแทน 3) เพราะบริษัทมีผลดำเนินงานดีจึงมี ROA สูง เพราะจะเป็นการลากเข้าความคือลากเข้าเป็น RM ทำให้ให้นักวิจัยใช้แต่ RM เสมอไป

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

Hair, Hult, Ringle และ Sarstedt (2017) แนะนำว่า

1) การวัดคุณภาพเครื่องมือกรณีที่เป็น FM ให้มุ่งเน้นที่ content validity, convergent validity, นัยสำคัญของน้ำหนัก (weights) และ multicollinearity และเนื่องจากตัวชี้วัดต้องเป็นอิสระต่อกันเหมือนใน MRA จึงใช้เครื่องมือเดียวกับ RM ไม่ได้

Convergent validity ใน FM ไม่ใช้การทดสอบแบบเดียวกับที่ใช้ใน RM การวิเคราะห์ convergent validity ให้ดูจากผลของ Redundancy analysis (ใช้ global item 1-3 ข้อ) หรือ Global item analysis (ใช้ global item 1 ข้อ)

2) ตัวชี้วัดใน RM สัมพันธ์กันสูงเนื่องจากเป็นสิ่งที่สะท้อนออกมาจากตัวแปรแฝงเดียวกัน เครื่องมือที่ใช้จะอาศัย content validity และสหสัมพันธ์คือ convergent validity (loading มีเครื่องหมายบวก มีนัยสำคัญ มีค่าไม่ต่ำกว่า 0.707) มี discriminant validity (ดูจาก Fornell-Larcker Criteria, Cross loading, HTMT), $AVE \geq 0.50$, $CR \geq 0.60$, Cronbach's Alpha ≥ 0.70

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

ผลเสียของการใช้ตัวแบบการวัดที่เป็น FM แต่นักวิจัยใช้เป็น RM

1. ทำให้การประมาณค่าสัมประสิทธิ์เส้นทางและพารามิเตอร์อื่นผิดพลาด (Diamantopoulos & Siguaw, 2006; MacKenzie et al., 2005)
2. ค่าสัมประสิทธิ์เส้นทางมามเอนเอียง (bias) อาจสูงเกินจริงหรือต่ำเกินจริง นำไปสู่การสรุปผลที่ผิดพลาด (Jarvis et al., 2003; Petter et al., 2007)
3. ใช้เกณฑ์ผิดประเภทในการประเมินความตรงและความเที่ยง (Bollen & Lennox, 1991; Diamantopoulos et al., 2008)
4. การตัดตัวชี้วัดทำให้แนวคิดเปลี่ยนไปและลดความตรงเชิงเนื้อหา (FM ห้ามตัดตัวชี้วัดแต่ RM สามารถตัดได้) (Diamantopoulos & Winklhofer, 2001; MacKenzie et al., 2011)
5. ประมาณค่าความคลาดเคลื่อนผิดพลาดเนื่องจากใน FM และ RM ความคลาดเคลื่อนอยู่ในตำแหน่งที่แตกต่างกัน (Edwards & Bagozzi, 2000)
6. จากการศึกษาเชิงประจักษ์ Jarvis et al. (2003) พบว่าการระบุตัวแบบผิดประเภทอาจทำให้ค่าพารามิเตอร์มามเอนเอียงสูงถึง 400%

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

ตัวอย่าง

Job Performance, Job Satisfaction, Stress, Leader-Member Exchange (LMX), Organizational Commitment เป็น RM

Organizational Culture, Service Quality, Experience Economy, Customer Recovery เป็น FM

Organizational Citizenship Behavior (OCB), Job Performance, Customer Loyalty, Employee Engagement, Service Quality, Innovation Capability, Perceived Value อาจเป็น RM หรือ FM

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

หมายเหตุ ตัวแปรแฝงเดียวกันสามารถวัดได้ทั้งในรูปแบบ FM หรือ RM ขึ้นอยู่กับว่าเรากำลังมองมันจากมุมมองไหน

1. FM เรามองจากองค์ประกอบ (sub-domain) หรือตัวชี้วัดร่วมกันสร้างหรือก่อตัวเป็นตัวแปรแฝง (Domain)
(องค์ประกอบ → ตัวแปรแฝง, ตัวชี้วัด → ตัวแปรแฝง)

2. RM เรามองจากผลการแสดงออกของตัวแปรแฝง (ตัวแปรแฝง → องค์ประกอบ, ตัวแปรแฝง → ตัวชี้วัด)

ยังมีตัวแบบผสม FM กับ RM เรียกว่าแบบจำลองผสม (Mixed Model) หรือ Multiple Indicator Multiple Cause Model (MIMIC Model)

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

หมายเหตุ เรามีเหตุผลที่ต้องใช้แบบจำลองผสม MIMIC Model หรือ model C ดังนี้

1. ต้องการความสมบูรณ์ของการวัด เพราะตัวแปรแฝงมีความซับซ้อนและมีหลายมิติ การวัดแบบผสมช่วยให้เข้าใจทั้งสาเหตุที่ก่อให้เกิดตัวแปรแฝงและผลที่แสดงออกมาจากตัวแปรแฝง
2. ต้องการลดข้อจำกัดของแต่ละแบบคือ FM มักมีปัญหาเรื่องค่าพารามิเตอร์เพราะรันด้วย MRA ที่ค่าของสัมประสิทธิ์ถดถอยเป็นไปได้ทั้งค่าบวก ค่าลบ มีนัยสำคัญ และไม่มีนัยสำคัญ ค่า weight ติดลบหรือไม่มีนัยสำคัญเป็นปัญหาสำคัญของการแปลผล ขณะที่ RM อาจไม่ครอบคลุมองค์ประกอบสำคัญบางประการ
3. เพิ่มความถูกต้องเชิงทฤษฎี ตัวแปรแฝงตามทฤษฎีบางตัวควรได้รับอิทธิพลจากตัวชี้วัดและควรส่งผลไปเป็นตัวชี้วัดไปในเวลาเดียวกัน การวิจัยแบบบูรณาการสามารถทดสอบทั้งปัจจัยที่เป็นสาเหตุและผลลัพธ์ในตัวแบบเดียวกัน
4. ลดความเอนเอียงในการวัดเพราะการใช้วิธีวัดหลายแบบช่วยตรวจสอบความตรงของการวัดตัวแปรแฝงเพราะในแบบจำลอง MIMIC ตัวแปรแฝงได้รับอิทธิพลจากตัวชี้วัด (formative indicators) และส่งผลต่อตัวชี้วัด (reflective indicators) จึงช่วยให้การวัดมีความสมบูรณ์และตรงตามทฤษฎีมากขึ้น โดยเฉพาะในกรณีของตัวแปรแฝงที่มีความซับซ้อน

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

ข้อที่ควรทราบและควรกังวลคือนักวิจัยทางสังคมศาสตร์เกือบทั้งหมดออกแบบสอบถามตาม sub-domain ที่มีมาตามทฤษฎี เช่น คุณภาพบริการมีไม่น้อยกว่า 5-7 sub-domain คือ tangibility, reliability, assurance, empathy, responsiveness หรือ tangibility, reliability, courtesy, competency, credibility, security, responsiveness และสร้างมาตรวัด sub-domain เหล่านี้โดยใส่ใจเรื่อง construct validity เป็นอย่างมาก แต่ตัวชี้วัดของตัวแปรแฝงนี้เป็น formative indicator model ซึ่งในขั้นตอนการวิเคราะห์ SEM ลูกศรจะชี้ออกจากตัวชี้วัดไปหาตัวแปรแฝง แต่นักวิจัยทุกคนหรือเกือบทุกคนใช้เป็น reflective model คือลูกศรชี้ออกจากตัวแปรแฝงไปหาตัวชี้วัด

การใช้ formative indicator มารัน SEM แบบ reflective indicator จะก่อให้เกิดความผิดพลาดคือ type I error สูง (H_0 : ตัวชี้วัดไม่สัมพันธ์กัน (จริง) แต่เอมารันแบบ RM และพบว่า loading สูง คือสหสัมพันธ์สูง ทำให้ปฏิเสธ H_0 นี่ก็คือ type I error) ผลการวิเคราะห์อาจผิดอย่างร้ายแรงแต่เราไม่รู้ตัวว่างานผิดพลาด ขณะเดียวกันการใช้ reflective indicator มารัน SEM แบบ formative indicator จะก่อให้เกิดความผิดพลาดแบบ type II error สูง ผลการวิเคราะห์อาจผิดพลาดอย่างร้ายแรงซึ่งเราก็ไม่รู้ตัวเช่นกัน

ความผิดพลาดส่วนใหญ่จะเป็นการใช้ formative indicator มารัน SEM แบบ reflective indicator

การประเมินคุณภาพโมเดลการวัดแบบ Reflective Measurement Model: RM

1. ความเที่ยง (Reliability)

Cronbach's alpha (ไม่น้อยกว่า 0.7)

Composite Reliability $CR \geq 0.7$, $AVE \geq 0.50$

2. ความตรงเชิงเหมือน (Convergent Validity ตาม SEM based threshold)

loadings > 0.707 แต่ถ้ามีค่าในช่วง 0.40-0.70 ให้คงไว้ได้ทั้งนี้ LV นั้นต้องมี $AVE \geq 0.50$, $CR \geq 0.6$ ต้องเป็นตัวชี้วัดสำคัญตรงตามวรรณกรรมถ้าตัดทิ้งไปจะทำให้ AVE และ CR จะลดลง

3. ความตรงเชิงจำแนก (Discriminant Validity)

ใช้ Fornell-Larcker criterion (ไม่แนะนำ) หรือ Cross-loadings หรือ Heterotrait-Monotrait Ratio (HTMT) (ไม่เกิน 0.85 หรือ 0.90)

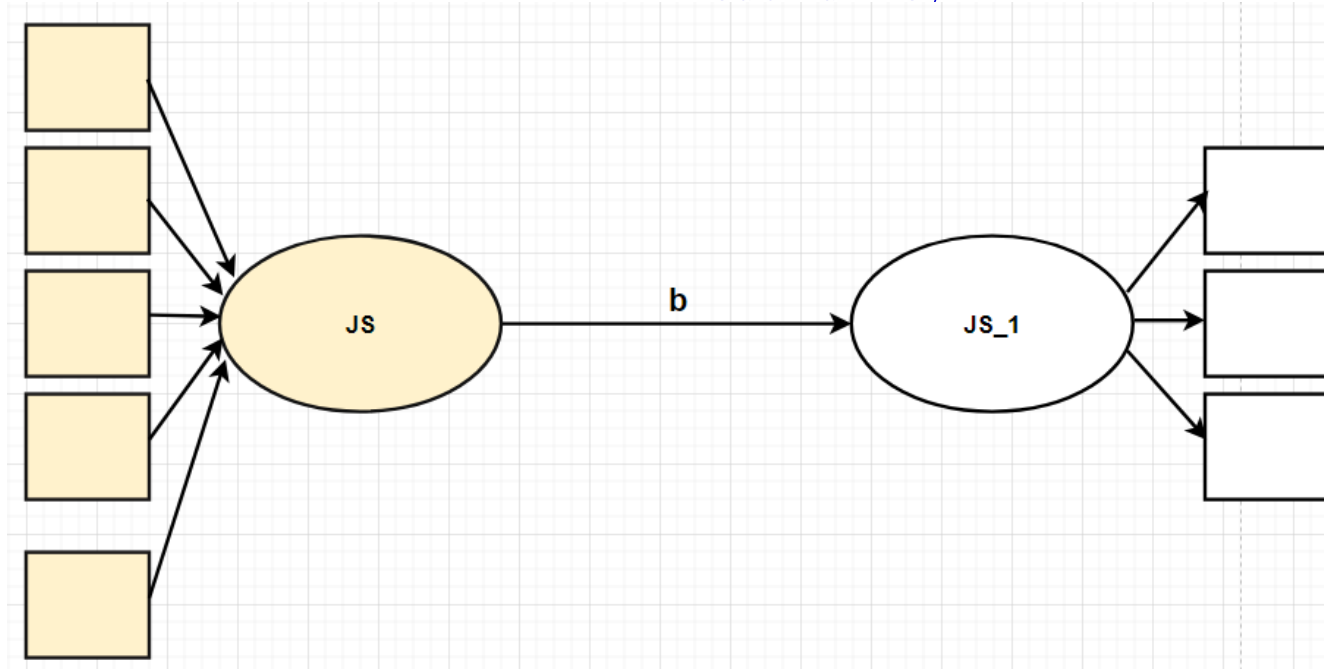
การประเมินคุณภาพโมเดลการวัดแบบ Formative Measurement Model: FM

การประเมินตัวแบบ FM ให้ประเมินที่ค่า Weights และ Multicollinearity

1. ตรวจสอบนัยสำคัญของ weights ซึ่งบ่งบอกความสำคัญของตัวชี้วัดแต่ละตัว
2. weights ที่ไม่มีนัยสำคัญอาจบ่งชี้ว่าตัวชี้วัดนั้นมีส่วนช่วยน้อยในการสร้างตัวแปรแฝง
3. weights สามารถเป็นค่าลบได้หากมีความสมเหตุสมผลทางทฤษฎี
4. ถ้า weight ไม่มีนัยสำคัญแต่ loading มีค่าไม่น้อยกว่า 0.707 หรือไม่น้อยกว่า 0.50 อาจพิจารณาคงตัวชี้วัดไว้ (คือให้พิจารณาจากทั้ง weight และ loading)
5. ไม่มีปัญหา Multicollinearity คือ $VIF \leq 5$ (หรือที่เข้มงวดกว่าคือ $VIF \leq 3$)
6. ทำการวิเคราะห์ Redundancy Analysis หรือ Global item approach หรือ Phantom item approach

การประเมินคุณภาพตัวแบบการวัดแบบ Formative Measurement Model: FM

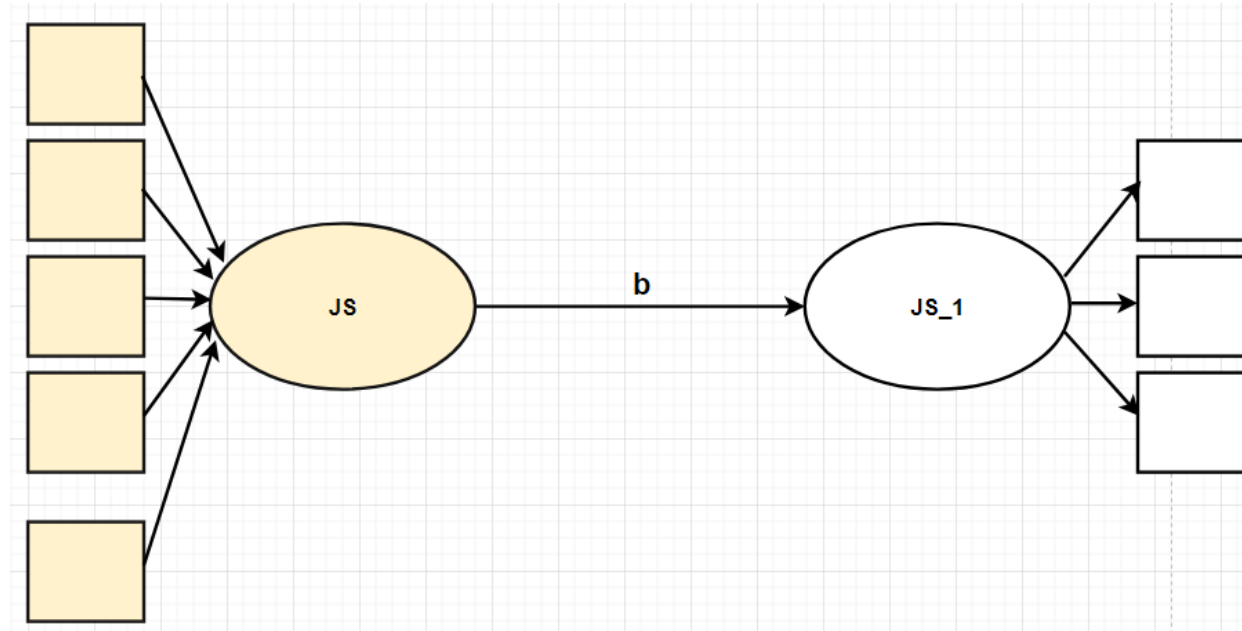
การวิเคราะห์ Redundancy



- 1) วางแผนแบบสอบถามให้มีตัวชี้วัดแบบภาพรวม (global items) สำหรับทำ redundancy analysis ไว้ล่วงหน้า 1-3 ข้อ เช่น ถามว่าในภาพรวมท่านพึงพอใจบริการที่นี้ เมื่อได้ข้อมูลมาแล้วให้ดำเนินการดังนี้
- 2) สร้างตัวแปรแฝงขึ้นมาตัวหนึ่ง อาจใช้ชื่อคล้ายกันกับในกรณี FM แต่วัดแบบ RM
- 3) โยงเส้นทางจากตัวแปรแฝงแบบ FM ไปยังตัวแปรแฝงแบบ RM ดังภาพ

การประเมินคุณภาพโมเดลการวัดแบบ Formative Measurement Model: FM

การวิเคราะห์ Redundancy



4) วิเคราะห์ SEM ตามกรอบข้างบนนี้ ถ้าค่า $b \geq 0.707$ (หรือ $R^2 \geq 0.5$) และมีนัยสำคัญทางสถิติแสดงว่าตัวชี้วัดแบบ FM มีความตรงเชิงเหมือน (ไม่ประเมิน discriminant validity เพราะต้องใช้สหสัมพันธ์ซึ่งไม่มีใน formative indicator) แปลว่าตัวชี้วัดแบบ FM สามารถวัดแนวคิดเดียวกันกับที่วัดด้วยวิธี RM ได้อย่างถูกต้อง ช่วยสร้างความมั่นใจว่า FM สามารถวัดแนวคิดได้ “อย่างแข็งแกร่ง ถูกต้อง ครบคลุม” เทียบเท่ากับการใช้ RM แต่มีความถูกต้องตรงบริบทการวัดกว่าการใช้เป็น RM โดยไม่คำนึงถึงธรรมชาติของ LV

การประเมินคุณภาพโมเดลการวัดแบบ Formative Measurement Model: FM

ตัวอย่าง Global Items สำหรับ Redundancy Analysis

ตัวอย่าง Global Items สำหรับ Customer Satisfaction

1. โดยรวมแล้ว ท่านมีความพึงพอใจในบริษัท/ผลิตภัณฑ์/บริการนี้ (1-5 หรือ 1-7 คะแนน)
2. ท่านจะให้คะแนนความพึงพอใจโดยรวมต่อบริษัท/ผลิตภัณฑ์/บริการนี้ในระดับใด (1-5 หรือ 1-7 คะแนน)
3. เมื่อพิจารณาทุกด้านของประสบการณ์การใช้บริการ/สินค้า ท่านรู้สึกพึงพอใจ (1-5 หรือ 1-7 คะแนน)

ตัวอย่าง Global Items สำหรับ Customer Loyalty

1. โดยรวมแล้ว กล่าวได้ว่าท่านมีความภักดีต่อบริษัท/แบรนด์นี้ (1-5 หรือ 1-7 คะแนน)
2. ท่านมีความภักดีต่อบริษัท/แบรนด์นี้ (1-5 หรือ 1-7 คะแนน)
3. เมื่อพิจารณาทุกด้านของประสบการณ์ ท่านยังคงเลือกที่จะใช้สินค้า/บริการของบริษัท/แบรนด์นี้ต่อไปในอนาคต (1-5 หรือ 1-7 คะแนน)

การประเมินคุณภาพโมเดลการวัดแบบ Formative Measurement Model: FM

ลักษณะสำคัญของ Global Items สำหรับ Redundancy Analysis

1. วัดแนวคิดโดยรวมของตัวแปรแฝง ไม่ลงรายละเอียดในแต่ละ sub-domain
2. ใช้คำที่แสดงภาพรวมของตัวแปรแฝงที่จะวัดแบบ RM เช่น

ความพึงพอใจโดยรวม

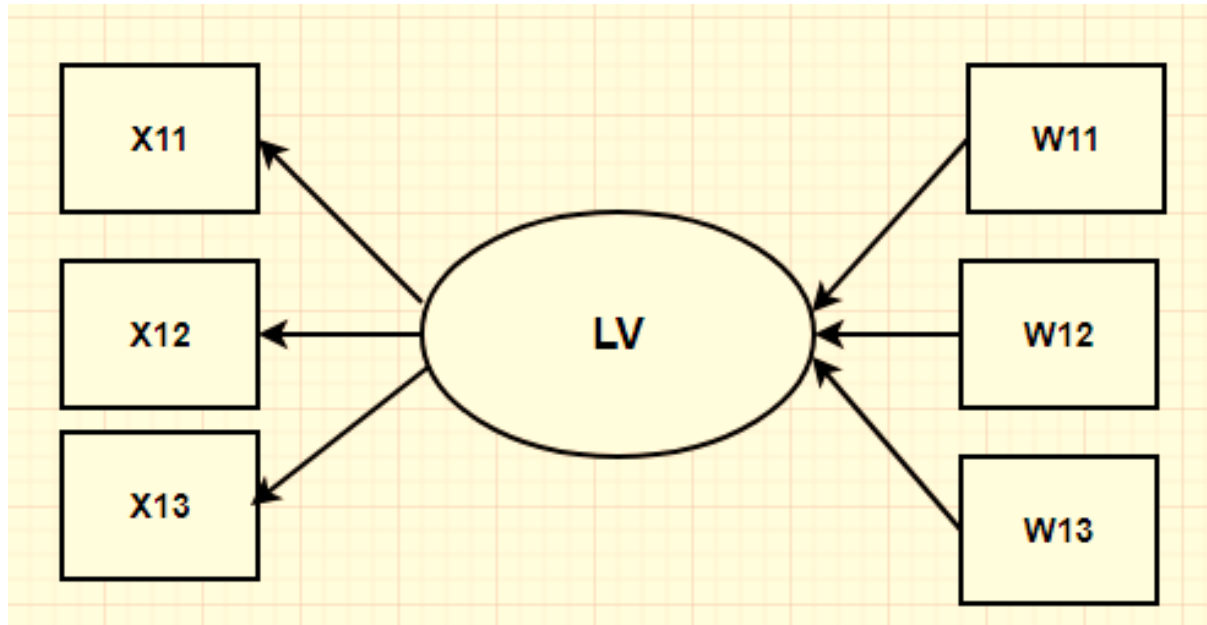
ความภาคภูมิใจโดยรวม

ความร่อกองค้การโดยรวม

3. มีจำนวนไม่มาก (มักใช้ 1-3 ข้อ) เพื่อไม่ให้แบบสอบถามยาวเกินไป

เมื่อเป็นดังนี้ก็แสดงว่าตัวแปรแฝงแบบ RM ที่เราสร้างมาล้อตัวแปรแฝงแบบ FM แต่วัดด้วย Global items มีคุณสมบัติของความเที่ยงตรงเชิงเหมือนและสามารถใช้เป็น reference ได้ และเมื่อผลการทดสอบพบว่า b มีค่าสูงและมีนัยสำคัญ (ใน SRA กรณี standardize model นั้น b เป็นได้ทั้งค่าสัมประสิทธิ์ถดถอยและสหสัมพันธ์) จึงแสดงว่า FM มีความสัมพันธ์อย่างมากกับ reference ที่เป็น RM ที่มีความเที่ยงตรงเชิงเหมือน FM จึงมีความเที่ยงตรงเชิงเหมือนด้วย

MIMIC measurement model (model C)



MIMIC (Multiple Indicator Multiple Causes model) เป็นตัวแบบการวัดที่ผสม FM กับ RM ในภาพข้างบน X11, X12, X13 คือตัวชี้วัดของตัวแบบ RM ส่วน W11, W12, W13 คือตัวชี้วัดของตัวแบบ FM การออกแบบผสมจะ ช่วยแก้ปัญหาที่ตัวแบบการวัดที่เป็นแบบ FM แต่นักวิจัยนำมาวิเคราะห์แบบ RM ซึ่งจะมีผลให้ผลการศึกษาผิดโดยที่ เราไม่รู้ตัวว่าผิด การแก้ไขคือสร้างมาตรวัดตามแบบ RM เพิ่มลงไปใน LV

การรันตัวแบบให้วัดตามนี้แล้วสั่งรันตามปกติแต่การวิเคราะห์คุณภาพเครื่องมือต้องมีการต่างกันดังนี้

การประเมินคุณภาพโมเดลการวัดแบบ MIMIC

1. ส่วน Formative Model (FM)

- 1) ตรวจสอบ weights ว่ามีนัยสำคัญทางสถิติหรือไม่
- 2) พิจารณาเครื่องหมายของ weights ว่าสอดคล้องกับทฤษฎีหรือไม่
- 3) ตรวจสอบปัญหา multicollinearity โดยดูจากค่า VIF (ควร ≤ 5 หรือที่เข้มงวดกว่าคือ ≤ 3)
- 4) หาก weight ไม่มีนัยสำคัญหรือมีเครื่องหมายไม่ตรงตามทฤษฎี ให้ตรวจสอบ loading เพิ่มเติม:
 - (1) ถ้า loading > 0.5 อาจพิจารณาคงตัวชี้วัดไว้
 - (2) ถ้า loading < 0.5 อาจพิจารณาตัดตัวชี้วัดออก แต่ต้องคำนึงถึงความตรงเชิงเนื้อหาด้วย

2. ส่วน Reflective Model (RM)

- 1) ตรวจสอบ loadings เพื่อวิเคราะห์ความตรงเชิงเหมือน
- 2) ตรวจสอบความเที่ยงด้วย Composite Reliability (CR), Cronbach's alpha, Cross loading, AVE

3. ความสัมพันธ์ระหว่างส่วน FM และ RM:

- 1) ดูค่า loading ของตัวชี้วัดในแบบ RM (ควรไม่ต่ำกว่า 0.707)
- 2) ดูค่า R^2 (ก็คือ loading_{ij}^2) ของตัวชี้วัดแบบ RM (ควรไม่ต่ำกว่า 0.5)

หากความสัมพันธ์ตรงเกณฑ์นี้แสดงว่าแบบ MIMIC มีความตรงเชิงเหมือน

Reflective Measurement Model (model A) vs Formative Measurement Model (model B)

1. ต่อไปนี้เป็นตัวอย่างตัวแปรแฝงและหัวข้อย่อยหรือองค์ประกอบ (sub-domain) ของตัวแปรแฝงนั้นๆ
คอลัมน์ซ้ายมือคือองค์ประกอบตามที่ปรากฏตามทฤษฎี ส่วนนี้ถือว่าเมื่อประกอบแล้วจะกลายเป็น
เรื่องราวของตัวแปรแฝงนั้น เรียกว่า composite model หรือ model B ถ้าวาดรูปใน SEM จะเห็นเป็นลูกศรชี้
จากองค์ประกอบหรือตัวชี้วัดเข้าหาตัวแปรแฝง

คอลัมน์ด้านขวา คือองค์ประกอบที่กำหนดขึ้นมาจากภายหลังเพื่อใช้เป็นแนวทางร่างแบบสอบถาม ตัวแปร
แฝงจะแสดงอาการออกมาเป็นตัวชี้วัด เรียกว่า reflective measurement model หรือ model A เมื่อวาดรูป
SEM จะเห็นเป็นรูปลูกศรชี้ออกจากตัวแปรแฝงไปยังตัวชี้วัด

2. การวิเคราะห์ข้อมูลซอฟต์แวร์จะรัน model B ด้วย MRA สถิติที่รายงานตัวแบบการวัดจะเป็น
สัมประสิทธิ์เส้นทาง เรียกว่า weight ค่า t ค่า p-value และ VIF และรัน model A ด้วย PLS algorithm สถิติที่
รายงานตัวแบบการวัดจะเป็น loadings, p-value, CR, alpha, AVE, HTMT

ตัวอย่าง sub-domain ใน Reflective model (model A) และ Formative model (model B)

Perceived Value

Formative Model (FM) Sub-domains

Quality Value

Price Value

Emotional Value

Social Value

Convenience Value

Functional Value

Epistemic Value

Conditional Value

Time/Effort

Risk Perception

Reflective Model (RM) Sub-domains

Overall Value Assessment

Satisfaction with Purchase

Repurchase Intention

Recommendation Willingness

Price Acceptability

Preference Over Alternatives

Usage Frequency

Positive Word-of-Mouth

Costs Switching Resistance

Willingness to Pay Premium

ตัวอย่าง sub-domain ใน Reflective model (model A) และ Formative model (model B)

Service Quality

Formative Model (FM) Sub-domains

Tangibility

Reliability

Responsiveness

Assurance

Empathy

Recovery

Flexibility

Access

Security

Professionalism

Reflective Model (RM) Sub-domains

Overall Quality Perception

Customer Satisfaction

Perceived Value

Behavioral Intentions

Word-of-Mouth Intent

Problem Resolution Satisfaction

Repurchase Intent

Brand Image Enhancement

Customer Loyalty

Complaint Behavior (inverse)

ตัวอย่าง sub-domain ใน Reflective model (model A) และ Formative model (model B)

Job Performance

Formative Model (FM) Sub-domains

Quality of Work
Quantity of Work
Meeting Deadlines
Error Rates (inverse)
Productivity Metrics
Goal Achievement
Efficiency
Innovation Output
Attendance
Supervisor Ratings

Reflective Model (RM) Sub-domains

Task Knowledge
Effort Expenditure
Interpersonal Facilitation
Job-Specific Skills
Time Management
Problem-Solving Ability
Adaptability
Technical Competence
Work Discipline
Initiative Taking

ตัวอย่าง sub-domain ใน Reflective model (model A) และ Formative model (model B)

Customer Loyalty

Formative Model (FM) Sub-domains

Repeat Purchase Behavior
Word-of-Mouth Recommendations
Share of Wallet
Price Premium Tolerance
Cross-Buying Behavior
Retention Duration
Brand Defense
Customer Referrals
Resistance to Competitor Offers
Forgiveness for Service Failures

Reflective Model (RM) Sub-domains

Service Quality
Product Quality
Perceived Value
Brand Image
Switching Costs
Trust
Commitment
Customer Experience
Loyalty Programs
Relationship Quality

ตัวอย่าง sub-domain ใน Reflective model (model A) และ Formative model (model B)

Organizational Justice

Formative Model (FM) Sub-domains

Distributive Justice
Procedural Justice
Interactional Justice
Informational Justice
Temporal Justice
Spatial Justice
Systemic Justice
Retributive Justice
Restorative Justice
Commutative Justice

Reflective Model (RM) Sub-domains

Trust in Management
Organizational Commitment
Job Satisfaction
Organizational Citizenship Behavior
Reduced Turnover Intention
Rule Compliance
Voice Behavior
Work Engagement
Job Performance
Reduced Counterproductive Behavior

ตัวอย่าง sub-domain ใน Reflective model (model A) และ Formative model (model B)

Destination Image

Formative Model (FM) Sub-domains

Natural Attractions

Cultural Heritage

Tourism Infrastructure

Accessibility

Local Cuisine

Hospitality

Safety and Security

Price/Value Perception

Entertainment Options

Accommodation Quality

Reflective Model (RM) Sub-domains

Destination Awareness

Visit Intention

Recommendation Willingness

Social Media Sharing

Length of Stay

Revisit Intention

Destination Preference

Positive Word-of-Mouth

Willingness to Pay Premium

Trip Planning Effort

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

การตรวจสอบคุณภาพเครื่องมือด้วยข้อมูลที่ได้รับ (measurement quality assessment)

1. ความเที่ยง (reliability) คือ ร้อยละของคะแนนจริงที่ปนอยู่ในค่าสังเกต หรือคือ สัดส่วนของคะแนนจริงที่มาตรวัดสามารถวัดได้ มาตรวัดที่มีความเที่ยงสูงคือ มาตรวัดที่มีสัดส่วนของความผันแปรของคะแนนจริงเทียบกับความผันแปรรวมสูง มีคุณสมบัติดังนี้

1) มีความคงเส้นคงวา (consistency) ข้อคำถามหรือส่วนต่าง ๆ ของเครื่องมือวัดนั้นวัดไปในทิศทางเดียวกัน คือ ทุกข้อควรสะท้อนแนวคิดเดียวกันของตัวแปรแฝง และมีความสัมพันธ์กันสูงคือ ถ้าตอบข้อหนึ่งในทิศทางหนึ่ง ควรตอบข้ออื่นในทิศทางเดียวกัน

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

2) ความเที่ยงเชิงโครงสร้าง (Equivalence) คือเมื่อนำไปใช้วัดสิ่งเดียวกันในสถานะหรือเงื่อนไขต่าง ๆ ก็จะได้ผลเหมือนเดิม เช่นไม่แปรไปตามรูปแบบของมาตรวัด เมื่อเป็นแบบสอบถามกระดาษหรือออนไลน์ ตอบเองหรือสัมภาษณ์ ไม่แปรตามสถานที่ตอบ ตอบที่บ้าน ตอบที่สำนักงาน หรือที่อื่นใด ไม่แปรตามอารมณ์หรือความวิตกกังวลของผู้ตอบ ไม่แปรตามสภาพแวดล้อมขณะตอบเช่นอุณหภูมิห้อง ไม่มีคำที่กำกวม ไม่เปิดช่องให้ผู้ตอบตอบเอาใจสังคม เช่น ต้องเชคที่หมายเลขสูง ๆ เพราะสังคมอยากเห็นแบบนั้น (social desirability) ไม่แปรตามลักษณะทางประชากรของผู้ตอบคือ เพศ อายุ สถานภาพการสมรส เชื้อชาติ ศาสนา การศึกษา ไม่แปรตามสังคมและวัฒนธรรมของผู้ตอบ ไม่แปรตามภาษา

3) มีความเสถียร (stability) คือเมื่อนำไปใช้ซ้ำในต่างกาลคือ เข้า บ่าย เย็น ต่างเดือน ต่างปี ก็ยังคงได้ผลเหมือนเดิม

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

วิธีวัดความเที่ยง

1. Cronbach's alpha คำนี้อาศัยวิธีวัดด้วย internal consistency คือวัดจากความสอดคล้องของตัวชี้วัดภายใน construct เดียวกัน "ความสอดคล้องภายในของเครื่องมือวัด (internal consistency)" หมายถึงการที่ส่วนต่าง ๆ ของเครื่องมือวัด (เช่น ข้อคำถาม) วัดไปในแนวทางเดียวกันหรือสอดคล้องกัน การวัดความสอดคล้องภายในช่วยให้มั่นใจว่าเครื่องมือวัดนั้นมีความเที่ยง (Reliability) สูง และสามารถวัดตัวแปรที่สนใจได้อย่างมีประสิทธิภาพ

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

Cronbach's alpha มีหลายสูตร

ก. Standard formula $\alpha = \frac{k}{k-1} \left(1 - \frac{\sum s_i^2}{s_t^2} \right)$ โดยที่ k = จำนวนตัวชี้วัดใน construct, s_i^2 = ความผันแปรของคะแนนของตัวชี้วัดที่ i , s_t^2 = ความผันแปรของคะแนนรวม

ข. Correlation-based formula $\alpha = \frac{k\bar{r}}{1+(k-1)\bar{r}}$ โดยที่ k = จำนวนตัวชี้วัดใน construct และ $\bar{r}$ คือค่าเฉลี่ยของ inter-item correlation ใน construct เดียวกัน สูตรนี้จะให้คำตอบคล้ายกับข้อ ก. ใช้ได้เมื่อตัวชี้วัดทุกตัวมีค่า s_i^2 เท่ากัน

ค. Split-Half Reliability Formula $\alpha = \frac{2r}{1+r}$ โดยที่ r คือค่าสหสัมพันธ์ระหว่างคะแนนครึ่งกลุ่ม

ง. Kuder-Richardson Formula 20 (KR-20) $\alpha = \frac{k}{k-1} \left(1 - \frac{\sum p_i q_i}{npq} \right)$, $p_i q_i$ = ความผันแปรของคะแนนของตัวชี้วัดที่ i , npq = ความผันแปรของคะแนนรวม

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

2. **Average corrected item-total correlation (ACITC)** คือค่าเฉลี่ยของสหสัมพันธ์ระหว่างตัวชี้วัดกับคะแนนรวมที่ปรับค่าของตัวชี้วัดตัวนั้น

3. **Average inter-item correlation (AIIC)** คือค่าเฉลี่ยของสหสัมพันธ์ระหว่างตัวชี้วัดภายใน construct เดียวกัน

ACITC และ AIIC เป็นวิธีการวัดความสอดคล้องภายใน (Internal Consistency)

CITC (corrected item-total correlation) วัดความสอดคล้องภายในว่าตัวชี้วัดแต่ละตัว (item) ขึ้นไปในทิศทางเดียวกันกับกลุ่ม (total) หรือไม่ ถ้าความสอดคล้องภายในสูงแสดงว่าเครื่องมือวัดให้ผลลัพธ์ที่สม่ำเสมอและคงที่ ให้ผลลัพธ์ที่เชื่อถือได้เมื่อใช้ซ้ำหรือในกลุ่มตัวอย่างต่างกัน

IIC (Inter-Item Correlation) วัดความสอดคล้องภายในคือวัดว่าตัวชี้วัดแต่ละตัว (Item) ขึ้นไปในทิศทางเดียวกันกับตัวชี้วัดอื่นหรือไม่ หากข้อคำถามมีความสอดคล้องกันสูง (ค่าสหสัมพันธ์สูง) แสดงว่าข้อคำถามเหล่านั้นวัดไปในทิศทางเดียวกัน

ถ้า CITC, IIC, ACITC, AIIC มีค่าสูงแสดงว่าเครื่องมือวัดให้ผลลัพธ์ที่สม่ำเสมอและคงที่ ให้ผลลัพธ์ที่เชื่อถือได้เมื่อใช้ซ้ำหรือในกลุ่มตัวอย่างต่างกัน

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

4. Composite reliability (CR) คือ Joreskog's rho_c, Dijkstra-Henseler's rho_a

$$\text{Rho}_c = \frac{(\sum_i^n \lambda_i)^2}{(\sum_i^n \lambda_i)^2 + \sum_i^n \delta_i^2}, \delta_i^2 = 1 - \lambda_i^2 = 1 - r_i^2 = \text{error variance, c หมายถึง composite}$$

$$\text{Rho}_a = \frac{(\sum_i^n w_i \lambda_i)^2}{(\sum_i^n w_i \lambda_i)^2 + \sum_i^n w_i^2 \delta_i^2}, a \text{ หมายถึง A ใน model A (reflective model)}$$

โดยที่ $w_i = \frac{\lambda_{ij}}{\sum_{j=1}^{m_i} \lambda_{ij}^2 + 2 \sum_{j < j'}^{m_i} \lambda_{ij} \lambda_{ij'}}$, λ_{ij} คือค่า loading ที่ j ของ construct ที่ i

หมายเหตุ 1. ถ้าว่า composite หมายถึง $\frac{\sum_{j=1}^{m_i} \lambda_{ij} x_{ij}}{\sum_{j=1}^{m_i} \lambda_{ij}}$ คือ construct score

2. Cronbach's alpha มักจะ under-estimate, Rho_c มักจะ over-estimate, Rho_a จะอยู่ระหว่างกลาง

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

5. AVE โดยที่ $AVE = \frac{\sum_i^n \lambda_i^2}{\sum_i^n \lambda_i^2 + \sum_i^n (1 - \lambda_i^2)} = \frac{\sum_i^n \lambda_i^2}{n}$ = ค่าเฉลี่ยของ item reliability

ค่าสหสัมพันธ์ λ_i ระหว่างค่าของตัวชี้วัดกับค่าของตัวแปรแฝง (λ_i คือ item loading หรือ simple regression coefficient) ต้องเป็นปริมาณบวก โดยตรรกะแล้วต้องเป็นบวกเท่านั้น เพราะตัวชี้วัดใช้วัดตัวแปรแฝงเดียวกัน สะท้อนจากสิ่งเดียวกัน ดังนั้นโดยนัยของ content validity ตัวชี้วัดจะต้องสัมพันธ์ทางบวกกับตัวแปรแฝงที่มุ่งวัด ถ้าพบว่ามีค่าลบให้แก้ไขโดยการ recode ข้อมูลก็กลับคะแนน 5→1, 4→2, 3→3, 2→4, 1→5 แล้วแก้ไขเนื้อหา/ภาษาของข้อถามให้กลับทางตามกันไป เช่น "ปวดศีรษะเนื่อง ๆ" → "ไม่ปวดศีรษะ"

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ 1

การวิเคราะห์ความกลมกลืน เฉพาะใน CB-SEM อาจเพิ่มการประเมิน CFA โดยมุ่งประเมินว่าตัวแบบกับข้อมูลสอดคล้องกันหรือไม่ นั่นคือ CFA ช่วยประเมินว่าข้อมูลนำเข้ากับข้อมูลนำออกเหมือนกันหรือไม่ ถ้าเหมือนกันก็แสดงว่าตัวแบบ good fit โดยการเปรียบเทียบ

อะเรย์ Σ คือ input correlation matrix จากไฟล์ข้อมูล กับ

อะเรย์ $\Sigma(\theta)$ คือ implied correlation matrix $\Sigma(\theta) = \Lambda\Phi\Lambda' + \Theta$ ที่ได้จากการวิเคราะห์ตัวแบบ

ว่าต่างกันหรือไม่ โดยที่ Λ = Factor loading matrix, Φ = Factor correlation matrix, Θ = Error variance matrix

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ 2

ในทฤษฎีการวัดเราทราบว่า

คะแนนจากการวัด (observe score) = คะแนนจริง (true score) + ความคลาดเคลื่อน (error)

หรือ $X = T + E$ ดังนั้น

ความเที่ยง = $\frac{V(T)}{V(X)}$ = หมายถึงสัดส่วนของคะแนนจริงที่ประเมินได้จากคะแนนสังเกต หรือ

ร้อยละของคะแนนจริงที่ปนอยู่ในค่าสังเกตหรือวัดได้ด้วยค่าสังเกต

ตัวอย่างเช่น ความเที่ยง = 0.80 หมายความว่า 80% ของคะแนนสังเกตเป็นคะแนนจริง ที่เหลืออีก 20% เป็นความคลาดเคลื่อน

ความเที่ยง = 1.00 หมายความว่า 100% ของคะแนนสังเกตเป็นคะแนนจริง ไม่มีความคลาดเคลื่อนปน

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

ความเที่ยง = $\frac{V(T)}{V(X)} = \frac{V(T)}{V(T+E)}$ เป็นการพิจารณาความผันแปรจากประชากรทั้งหมด ไม่ใช่จากข้อมูล

เดียว ๆ ของคนใดคนหนึ่งคือไม่ใช่ $\frac{T}{T+E}$ โดยที่

$V(T)$ คือความผันแปรของคะแนนจริง (true score variance) แสดงความแตกต่างที่แท้จริงระหว่างบุคคล

$V(T+E)$ คือความผันแปรทั้งหมดที่วัดได้ (observed score variance) รวมความแตกต่างที่แท้จริงและความคลาดเคลื่อน

อัตราส่วนนี้บอกว่า "ความผันแปรของคะแนนที่วัดได้มีส่วนที่เป็นความผันแปรจริงเท่าไร" หรือ "ข้อมูลมีความแม่นยำมากน้อยเพียงใด"

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

E (ความคลาดเคลื่อนในการวัด) ประกอบด้วย

1. ความคลาดเคลื่อนแบบสุ่ม (Random Error) ประกอบด้วย 1) ความเหนื่อยล้าของผู้ตอบ 2) สภาพแวดล้อมในการทดสอบ (เสียงรบกวน, อุณหภูมิ) 3) ความผันผวนของสภาวะทางอารมณ์ 4) การเดาคำตอบ 5) ความคลาดเคลื่อนในการให้คะแนน
2. ความคลาดเคลื่อนเชิงระบบ (Systematic Error) ประกอบด้วย 1) ความลำเอียงในการตอบ (Response bias) 2) ความโน้มเอียงที่จะเห็นด้วย (Acquiescence bias/yea-saying bias) เพราะไม่อยากขัดแย้ง ขี้เกียจคิด เป็นคนสุภาพ ไม่แน่ใจคำถาม คำถามยาวหรือยาก 3) ความต้องการเป็นที่ยอมรับทางสังคม (Social desirability) คือแสดงพฤติกรรมเกินจริง ตอบตามที่สังคมยอมรับ ตอบตามที่สังคมคาดหวัง 4) การตอบแบบทิศทางเดียว (Response set) เช่น ไม่แน่ใจ เห็นด้วย เห็นด้วยอย่างยิ่ง ซ้ำๆ 5) ความลำเอียงของผู้ประเมิน (Rater bias) 6) ความไม่สม่ำเสมอในความยากของข้อคำถาม
3. ความคลาดเคลื่อนเฉพาะบุคคล (Person-specific error) ประกอบด้วย 1) ความเข้าใจในข้อคำถามที่แตกต่างกัน 2) ทักษะทางภาษาที่แตกต่างกัน 3) แรงจูงใจในการทำแบบทดสอบ

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

2. ความตรง (Validity) คือปริมาณที่ชี้ให้เห็นความแม่นยำ (accuracy) ของมาตรวัดว่าสามารถวัดสิ่งที่ต้องการวัดได้ตรงตามประสงค์ แยกเป็นตรงเนื้อหา และตรงไปที่ตัวแปรแฝงเดียวกันไม่ใช่วัดตัวแปรแฝงอื่น วิธีวัดมีดังนี้

- 1) Content validity มาตรวัดสามารถวัดเนื้อหาได้ครบถ้วนและตรงเนื้อหา ประเมินได้โดยอาศัยผู้เชี่ยวชาญเครื่องมือ เรียกการประเมินว่า Item-Objective Congruency (IOC) คำว่า Objective ก็คือตัวแปรแฝงที่ต้องการจะวัด
- 2) Construct validity คือสัดส่วนที่มาตรวัดสามารถวัดตัวแปรแฝงได้ถูกต้องตามทฤษฎีของตัวแปรนั้น แยกเป็น

(1) convergent validity ตัวชี้วัดต้องสัมพันธ์กันเองเพราะวัดถึงเรื่องของตัวแปรแฝงเดียวกันและสัมพันธ์อย่างมากกับตัวแปรแฝง คือวัดตรงตัวแปรแฝงเพราะเป็นภาพสะท้อนจากตัวแปรแฝง

หมายเหตุ ในศาสตร์การวัดผล convergent validity คือผลการวัดสหสัมพันธ์ของคะแนนจากแบบทดสอบ 2 ชุดขึ้นไป ถ้าพบว่าสหสัมพันธ์ระหว่างคะแนนชุดต่าง ๆ มีค่าสูงก็แสดงว่ามีความตรงเชิงเหมือน แต่ในกรณีที่เรามีแบบสอบถามชุดเดียวจึงใช้วิธีนำคะแนนตัวชี้วัดกับ construct score มาหาค่าสหสัมพันธ์เรียกว่า item loading ถ้าพบว่ามีค่าเกิน 0.707 เป็นบวก และมีนัยสำคัญ $AVE \geq 0.50$ แสดงว่ามีความตรงเชิงเหมือน คือเหมือนกับเรื่องราวของตัวแปรแฝงที่กำลังวัด น่าจะเรียกว่า Uni-dimensionality = ความมั่นใจว่าเรากำลังวัดสิ่งเดียวกันไม่ใช่หลายสิ่งปนกัน

- (2) Discriminant validity ตัวชี้วัดของตัวแปรแฝงหนึ่งไม่สัมพันธ์สูงกับตัวชี้วัดของตัวแปรแฝงอื่น

การประเมินความตรงในแบบ Second-order model

แบบ Second-order model จะมี Domain หรือ construct หรือ Second-order construct กับ Sub-domain หรือ Sub-construct หรือ First-order construct เราสามารถวิเคราะห์ inner model ได้ 2 วิธี

1. Two-step approach ใช้วิธีสร้างค่าให้แก่ sub-domain อาจเป็นค่าเฉลี่ยหรือคะแนนปัจจัย แล้วใช้ค่าเหล่านี้ในฐานะตัวชี้วัด วิธีนี้ค่า loading ไม่ใช่ indicator loading ห้ามใช้ประเมินความตรงแต่ใช้แสดงความสัมพันธ์ระหว่าง domain กับ sub-domain การประเมินความตรงให้แยกไปทำต่างหาก

2. Repeated indicator approach ให้จัดตัวชี้วัดให้แก่ sub-domain ก่อนแล้วใช้ค่าตัวชี้วัดนั้นโดยจัดให้เป็นตัวชี้วัดของ domain ด้วยวิธีนี้ค่า loadings ของ domain จึงใช้เป็น item loading ค่าจะไม่เท่ากับ loading ใน sub-domain ห้ามนำ loading ของ sub-domain มาใช้เป็น loading

3. ใน CFA ให้ใช้ทุก sub-domain มาวิเคราะห์พร้อมกัน สิ่งที่ได้คือผลการทดสอบความกลมกลืนคือทดสอบว่า input correlation matrix ของตัวชี้วัด กับ implied correlation matrix ของตัวชี้วัดไม่ต่างกันเพื่อแสดง model-data fit และยังสามารถใช้ค่า loading ไปประเมินคุณภาพมาตรวัดคือ convergent validity, HTMT, AVE, CR, Alpha, Fornell-Larcker criteria

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

เราสามารถวัด Discriminant validity ได้ 3 วิธี

วิธีที่ 1 Heterotrait-Monotrait Ratio (HTMT) (Henseler, Ringle, and Sarstedt, 2015) วิธีนี้เราจะคำนวณหาค่าสหสัมพันธ์ไขว้ระหว่างตัวชี้วัดของ construct คู่ใด ๆ ให้อยู่ ๆ ทำทีละคู่จนครบทุกคู่ แล้วเฉลี่ยเรียกว่า Heterotrait correlation จากนั้นหาค่าเฉลี่ยสหสัมพันธ์ของตัวชี้วัดในแต่ละตัวแปรแฝงแล้วนำมาหาค่าเฉลี่ยแบบ Geometric mean เรียกว่า Monotrait correlation แล้วหารกัน เกณฑ์ตัดสินใจคือ $HTMT_{ij} \leq 0.85, 0.90$

$$HTMT_{ij} = \frac{1}{K_i K_j} \sum_{g=1}^{K_i} \sum_{h=1}^{K_j} r_{ig,jh} \div \left(\frac{2}{K_i (K_i - 1)} \sum_{g=1}^{K_i-1} \sum_{h=g+1}^{K_i} r_{ig,ih} \frac{2}{K_j (K_j - 1)} \sum_{g=1}^{K_j-1} \sum_{h=g+1}^{K_j} r_{jg,jh} \right)^{\frac{1}{2}}$$

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

วิธีที่ 2 Cross loading คือค่าสหสัมพันธ์ระหว่างคะแนนตัวชี้วัดกับคะแนนปัจจัย (construct score) โดยเราจะหาค่าสหสัมพันธ์ระหว่างทุกตัวชี้วัดกับทุกตัวแปรแฝง ค่าสหสัมพันธ์ของตัวแปรแฝงกับตัวชี้วัดของตนต้องสูงกว่ากับตัวชี้วัดของตัวแปรแฝงอื่น

ค่า cross loading ยังช่วยตรวจสอบและแก้ไขปัญหา Discriminant validity โดยวิธี HTMT ได้ด้วย กล่าวคือ ถ้าหากพบว่า HTMT ของตัวแปรแฝงคู่ใดมีค่าสูงเกิน 0.90 หรืออาจเกิน 1.00 แสดงว่าตัวแปรแฝงคู่นั้นมีตัวชี้วัดที่ซ้ำกันให้ตรวจสอบจากตาราง cross loading ว่าสาเหตุเกิดจากตัวชี้วัดใดแล้วตัดตัวชี้วัดที่สัมพันธ์กับตัวแปรอื่นสูงเกิน 0.707 ทิ้ง

วิธีที่ 3 Fornell-Larcker Criteria (Fornell and Larcker, 1981) วิธีนี้หาค่าสหสัมพันธ์ระหว่างตัวแปรแฝงโดยใช้ Construct score แต่สหสัมพันธ์ในตนเอง (ซึ่งควรจะเท่ากับ 1) กลับใช้วิธีถอดรากที่ 2 ของ AVE_i โดยที่

$$AVE_i = \frac{1}{k} \sum_{j=1}^k \lambda_{ij}^2 \text{ ทั้งนี้ } \lambda_{ij}^2 \text{ คือ (item correlation)}^2$$

เกณฑ์ตัดสินใจก็คือค่าในแนวทแยงของเมทริกซ์สหสัมพันธ์นี้ต้องสูงกว่าค่าสหสัมพันธ์ที่ตัวแปรเป้าหมายนั้นมีกับตัวแปรอื่น

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ 1. สูตรใหม่ของ HTMT คือ HTMT2 สูตรนี้ใช้ Geometric mean ทั้งหมด ขณะนี้มีใช้ใน ADANCO

$$\text{HTMT2}_{ij} = \frac{\overbrace{\sqrt{K_i \cdot K_j \prod_{g=1}^{K_i} \prod_{h=1}^{K_j} r_{ig,jh}}}^{\text{geometric mean of indicator correlations between } \xi_i \text{ and } \xi_j}}{\sqrt{\underbrace{\sqrt{\frac{K_i^2 - K_i}{2} \prod_{g=1}^{K_i-1} \prod_{h=g+1}^{K_j} r_{jg,jh}}}_{\text{geometric mean of indicator correlations within } \xi_i} \cdot \underbrace{\sqrt{\frac{K_j^2 - K_j}{2} \prod_{g=1}^{K_j-1} \prod_{h=g+1}^{K_i} r_{ig,jh}}}_{\text{geometric mean of indicator correlations within } \xi_j}}}$$

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ 2 HTMT ดีกว่า Fornell-Larcker Criteria ตรงที่

HTMT ใช้สารสนเทศจากสหสัมพันธ์แท้ ๆ คือสหสัมพันธ์ระหว่างตัวชี้วัด ขณะที่ Fornell-Larcker Criteria ใช้สหสัมพันธ์ระหว่าง Factor score ซึ่งปิดบังค่าสหสัมพันธ์แท้ ๆ อันเป็นรากฐานเอาไว้

Fornell-Larcker Criteria หาค่าสหสัมพันธ์ในตน (self correlation) แบบอ้อม ๆ คืออาศัยรากที่ 2 ของ AVE_i

เราทราบว่า $AVE_i = \frac{1}{k} \sum_{j=1}^k \lambda_{ij}^2$ คือค่าเฉลี่ยของ item reliability และ item reliability คือกำลังสองของ item correlation การใช้ $\sqrt{AVE_i}$ เป็นค่าสหสัมพันธ์ในตนจึงไม่น่าจะถูกต้องเพราะใช้แบบอ้อม ๆ

Fornell-Larcker Criteria ประเมินปัญหา discriminant validity ผิดถ้า construct คู่ใด ๆ มีความคล้ายคลึงกันมากทำให้ construct correlation มีค่าสูง เช่นมีค่าสหสัมพันธ์ระหว่าง construct score แต่ละเกณฑ์คือ 0.85 กรณีนี้ Fornell-Larcker Criteria จะอนุญาตถ้ารากที่ 2 ของ AVE_i มีค่ามากกว่า

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ 3 HTMT และ HTMT2 ใช้ประเมิน Discriminant validity ได้ไม่แตกต่างกันมาก แต่มีการพัฒนา HTMT2 เพื่อปรับปรุงข้อบกพร่องที่ใช้ใน HTMT

ในการพัฒนา HTMT ตกลงว่า loading ต้องมีค่าเท่ากันแต่ค่าเฉลี่ยของตัวชี้วัดจากคะแนน 1 ถึง 5 สามารถแปรไปได้ เรียกว่า Tau-Equivalent ซึ่งไม่จริงในทางปฏิบัติ

แต่การพัฒนา HTMT2 ถือว่า loading และค่าเฉลี่ยสามารถแปรค่าไปได้ตรงตามที่จะเป็นในทางปฏิบัติ เรียกว่า Congeneric

HTMT2 เปลี่ยนไปใช้ geometric mean (GM) คือ $\sqrt[k]{r_{11} * r_{12} * \dots * r_{1k}}$ แทนค่าเฉลี่ยเลขคณิต (AM) เพราะ $\frac{r_{11} + r_{12} + \dots + r_{1k}}{k}$ จะเปลี่ยนแปลงไปมากถ้า r_{ij} มีค่าผันแปรมากขณะที่ $\sqrt[k]{r_{11} * r_{12} * \dots * r_{1k}}$ ไม่เปลี่ยนแปลงมาก

กิจกรรมที่ต้องดำเนินการในการวิเคราะห์ SEM

หมายเหตุ 4

$AM = \frac{X_1 + X_2 + \dots + X_n}{n}$ ใช้ในสถานการณ์ทั่วไปที่ข้อมูลไม่ผันผวนวูบวาบ

$GM = \sqrt[n]{X_1 * X_2 * \dots * X_n}$ ใช้ในสถานการณ์ที่

1. ข้อมูลมีการเปลี่ยนแปลงแบบทวีคูณ (Multiplicative Growth) เช่น อัตราการเติบโตทางเศรษฐกิจ ผลตอบแทนการลงทุน การเติบโตของประชากร การขยายตัวทางเทคโนโลยี การแพร่กระจายของโรค
 2. ข้อมูลที่มีการกระจายแบบไม่สมมาตร (skewed distribution) เช่น ราคาหุ้น รายได้ ราคาที่ดิน
 3. ข้อมูลที่มีช่วงค่าแตกต่างกันมาก (wide range) เช่น ข้อมูลทางการเงิน ข้อมูลราคาสินค้า
 4. ข้อมูลที่เป็นอัตราส่วนหรือร้อยละ (ratios or percentages) เช่น อัตราส่วนทางการเงิน ดัชนีราคา
- อัตราการเปลี่ยนแปลง

เมื่อนำมาหาค่าเฉลี่ย GM จะน้อยกว่า AM

เกณฑ์ตัดสินใจความเที่ยง

1. Cronbach's Alpha ≥ 0.70 (Nunnally & Bernstein, 1994)

ถ้า Cronbach's Alpha = 0.70 หมายความว่ามาตรวัดสามารถวัดคะแนนจริงได้ 70% หรือ 70% ของคะแนน
สังเกตคือค่าคะแนนจริง/ความจริงของมโนทัศน์นั้น ที่ขาดไป 30% เรียกว่ามีความคลาดเคลื่อนจากการวัด คือ

1) random error เป็นความคลาดเคลื่อนที่แตกต่างกันไปในแต่ละบุคคล เช่น อารมณ์เสีย ตอบโดยไม่อ่าน
คำถาม ตอบคำปานกลางเสมอ ๆ รีบตอบให้เสร็จ ๆ ไป งงกับคำถามและตอบไปแบบงง ๆ ไม่เข้าใจคำถาม
ตีความคำถามแตกต่างกันไปตามเวลา คือเว้นไปครู่หนึ่งกลับมาอ่านคำถามใหม่กลับตีความเป็นอย่างอื่น

2) Systematic error เป็นความคลาดเคลื่อนที่กระทบกับทุกคน

(1) item problem คำชี้แจงให้ตอบคำถามไม่ชัดเจน ใช้ภาษาไม่ดี คำกวม ประโยคซับซ้อน แปลมาผิด
ข้อถามวัดเรื่องอื่น เป็นความคลาดเคลื่อนที่ทำให้ผู้ตอบไม่เข้าใจคำถาม

(2) Situational factor ช่วงเวลาของวัน สิ่งแวดล้อมขณะตอบ พลังที่ต้องใช้ตอบ บริบทที่ต่างออกไป
ด้วยเหตุนี้เราจึงต้องออกแบบคำถามที่ข้อความต้องกระชับเพื่อให้ผู้ตอบล้า ออกแบบไว้มากขึ้นตาม content
validity แล้วค่อย ๆ คัดกรองออก

เกณฑ์ตัดสินใจความเที่ยง

2. Composite Reliability (CR)

Joreskog's rho-c ≥ 0.60 (Fornell, & Larcker, 1981)

Dijkstra-Henseler's rho-a ≥ 0.60 (Dijkstra, & Henseler, 2015)

3. Average Variance Extracted

$AVE_i \geq 0.50$ (Fornell & Larcker, 1981)

เกณฑ์ตัดสินใจความเที่ยง

4. Average Corrected item-total correlation (ACITC) (Field, 2013).

ACITC > 0.50 Acceptable ใช้วัดภาพรวมของมาตรวัด ส่วน CITC ใช้ประเมินรายข้อดังนี้

CITC < 0.15: Very poor, strong candidate for removal น่าจะวัดคนละเรื่อง

0.15 - 0.19: Poor, consider removing

0.20 - 0.29: Fair, possible candidate for removal

0.30 - 0.39: Good

0.40 - 0.49: Very good

0.50 - 0.59: Excellent

0.60 - 0.69: Very excellent

0.70 - 0.79: Outstanding

0.80 - 0.89: Very high, possible item redundancy

CITC $\geq$ 0.90: Extremely high, strong indication of item redundancy

สรุปว่า CITC < 0.30 poor/unacceptable

0.30 – 0.70 Good/Acceptable

> 0.70 very good but check for redundancy

เกณฑ์ตัดสินใจความเที่ยง IIC

5. Average inter-item correlation, AIIC

$AIIC \leq 0.20$ Low, $AIIC = 0.20$ to 0.36 Medium, $AIIC \geq 0.40$ High (Cohen, 1988)

1) $IIC = 0.15$ to 0.50 General Acceptable Range (Clark & Watson, 1995)

2) $IIC = 0.20$ to 0.40 Optimal/Recommended Range (Briggs & Cheek, 1986)

3) ถ้า $IIC \geq 0.40, 0.50$ อาจมีข้อถามอาจซ้ำกัน

พิสูจน์สูตร CR, AVE และ correlation-based alpha

พิสูจน์สูตร composite reliability (CR)

$$Y = \text{ยอดรวมคะแนนสังเกต} = (\lambda_1\xi + \delta_1) + (\lambda_2\xi + \delta_2) + \cdots + (\lambda_k\xi + \delta_k)$$

$$= \lambda_1\xi + \lambda_2\xi + \cdots + \lambda_k\xi + (\delta_1 + \delta_2 + \cdots + \delta_k) = T + E$$

$$V(Y) = (\lambda_1 + \lambda_2 + \cdots + \lambda_k)^2 V(\xi) + V(\delta_1 + \delta_2 + \cdots + \delta_k)$$

$$= (\sum_{i=1}^k \lambda_i)^2 + \sum_{i=1}^k V(\delta_i)$$

$$V(Y) = (\sum_{i=1}^k \lambda_i)^2 + \sum_{i=1}^k (1 - \lambda_i^2) = V(T) + V(E) \text{ ดังนั้น}$$

$$CR = \frac{V(\text{true score})}{V(\text{observed score})} = \frac{V(T)}{V(Y)} = \frac{(\sum_{i=1}^k \lambda_i)^2}{(\sum_{i=1}^k \lambda_i)^2 + \sum_{i=1}^k (1 - \lambda_i^2)}$$

พิสูจน์สูตร Average Variance Extracted: AVE

จากสมการถดถอย (สมการการวัด) $Y_i = \lambda_i \xi + \delta_i ; i = 1, 2, \dots, k$

Coefficient of determination คือสัดส่วนความผันแปรที่ตัวแปรอิสระอธิบายความผันแปรรวมได้คือ

$$r_i^2 = \frac{\sum_i^k \hat{y}_i^2}{\sum_i^k y_i^2} = \frac{SS(\text{Regression})}{SS(\text{Total})} = \frac{V(\lambda_i \xi)}{V(Y_i)} = \frac{\lambda_i^2 V(\xi)}{\lambda_i^2 V(\xi) + V(\delta_i)} = \frac{\lambda_i^2}{\lambda_i^2 + (1 - \lambda_i^2)}$$

เมื่อเฉลี่ยจะได้ค่าเฉลี่ยของ r_i^2 เรียกว่า Average Variance Extracted (AVE)

$$AVE_i = \frac{1}{k} \sum_i^k \frac{\lambda_i^2}{\lambda_i^2 + (1 - \lambda_i^2)} = \frac{\frac{1}{k} \sum_i^k \lambda_i^2}{\frac{1}{k} [\sum_i^k \lambda_i^2 + \sum_i^k (1 - \lambda_i^2)]} = \frac{1}{k} \sum_i^k \lambda_i^2$$

$$\text{หรือ } AVE_i = \frac{\sum_i^k \lambda_i^2}{\sum_i^k \lambda_i^2 + \sum_i^k (1 - \lambda_i^2)}$$

การพิสูจน์ว่า Cronbach's alpha คือ $\alpha = \frac{k}{k-1} \left(1 - \frac{\sum_i^k \sigma_i^2}{\sigma_t^2} \right) = \frac{k\bar{r}}{1+(k-1)\bar{r}}$

$$\text{จาก } \alpha = \frac{k}{k-1} \left(1 - \frac{\sum_i^k \sigma_i^2}{\sigma_t^2} \right)$$

พิจารณาที่ $\frac{\sum_i^k \sigma_i^2}{\sigma_t^2}$ จะพบว่า

$$\begin{aligned} 1. \quad \sigma_t^2 &= V(X_1 + X_2 + \dots + X_k) \\ \sigma_t^2 &= \sum_i^k V(X_i) + 2 \sum_{i < j} V(X_i X_j) \\ &= \sum_i^k \sigma_i^2 + 2 \sum_{i < j} \sigma_{ij} \end{aligned}$$

$$\text{และเราทราบว่า } r_{ij} = \frac{\sigma_{ij}}{\sigma_i \sigma_j} \text{ ดังนั้น } \sigma_{ij} = r_{ij} \sigma_i \sigma_j$$

และเนื่องจาก all possible correlation ของ $X_1, X_2, \dots, X_k$ มีได้ $C(k,2)$ คู่ คือ $\frac{k(k-1)}{2}$ คู่

$$\text{ดังนั้น } \bar{r} = \frac{1}{\frac{k(k-1)}{2}} \sum \sum_{i < j} r_{ij} \text{ นั่นคือ } \sum \sum_{i < j} r_{ij} = \frac{k(k-1)}{2} \bar{r}$$

2. กำหนดให้ σ_i^2 คงที่เท่ากับ a ดังนั้น

$$\sum_i^k \sigma_i^2 = \sum_i^k a = ka$$

$$\begin{aligned} \text{และ } \sigma_t^2 &= \sum_i^k \sigma_i^2 + 2 \sum_{i < j} \sigma_{ij} = ka + 2a \sum_{i < j} r_{ij} = ka + 2a \frac{k(k-1)}{2} \bar{r} \\ &= ka + 2a \frac{k(k-1)}{2} \bar{r} = ka[1 + (k-1)\bar{r}] \end{aligned}$$

$$\begin{aligned} \text{ดังนั้น } \alpha &= \frac{k}{k-1} \left(1 - \frac{\sum_i^k \sigma_i^2}{\sigma_t^2} \right) = \frac{k}{k-1} \left(1 - \frac{ka}{ka[1+(k-1)\bar{r}]} \right) \\ &= \frac{k}{k-1} \left(1 - \frac{1}{[1+(k-1)\bar{r}]} \right) = \frac{k}{k-1} \left(\frac{(k-1)\bar{r}}{1+(k-1)\bar{r}} \right) \\ &= \frac{k\bar{r}}{1+(k-1)\bar{r}} \end{aligned}$$

เกณฑ์ตัดสินใจความตรง (validity)

1. $IOC \geq 0.50$ (content validity) และมีนัยสำคัญ (Rovinelli & Hambleton, 1977; Turner, & Carlson, 2003)
2. Item loading ≥ 0.707 มีนัยสำคัญ และเป็นปริมาณบวก ถ้ามีค่าต่ำกว่า 0.707 ต้องมีค่าไม่ต่ำกว่า 0.50 พร้อมกันกับที่ $AVE \geq 0.50$ (convergent validity) และมีนัยสำคัญ (Chin, 1998) เกณฑ์ปกติจะกำหนดให้ loading มีค่าเป็นบวกที่ไม่ต่ำกว่า 0.707 และมีนัยสำคัญ ถ้า loading < 0.40 ให้ตัดตัวชี้วัดนั้นทิ้ง

แต่ก็ผ่อนปรนยอมให้คงตัวชี้วัดที่มีค่า loading ต่ำแต่อยู่ในช่วงระหว่าง 0.40-0.70 ได้เพราะจะทำให้จำนวนตัวชี้วัดเหลือมากพอเหมือนเดิม แต่ก็ต้องสอดคล้องกับเงื่อนไขดังต่อไปนี้ (Hair et al., 2018)

- 1) loading ค่าอื่นต้องไม่ต่ำกว่า 0.707
 - 2) $CR \geq 0.7$
 - 3) $AVE \geq 0.5$
 - 4) ถ้าทิ้งตัวชี้วัดนั้นไปจะทำให้ความตรงลดลง
 - 5) เป็นตัวชี้วัดสำคัญทางทฤษฎีคือมีความตรงเชิงเนื้อหา ตรงทฤษฎี ตรงเชิงเหมือน ตรงเชิงจำแนก
3. $HTMT \leq 0.85, 0.90$ และมีนัยสำคัญ $H_0: HTMT \geq 0.85$ (หรือ 0.90) vs $H_1: HTMT < 0.85$ (หรือ 0.90) (Henseler, Ringle & Sarstedt, 2015)

เกณฑ์ตัดสินใจสำหรับ R-square

R-square คือสัดส่วนหรือร้อยละของความผันแปรในตัวแปรเป้าหมายที่ตัวแปรสาเหตุสามารถควบคุมหรืออธิบายได้ (proportion explained) ถ้ามีค่าสูงแสดงว่าตัวแปรสาเหตุสามารถอธิบายได้ว่าตัวแปรผลลัพธ์ส่วนใหญ่ผันแปรไปเพราะปัจจัยสาเหตุเหล่านั้น

0.02: Small effect, 0.13: Medium effect 0.26: Large effect (Cohen,1988).

0.25: Weak, 0.50: Moderate, 0.75: Substantial (Hair et al.,2011)

0.10: Minimum acceptable level (Falk and Miller, 1992)

0.19: Weak, 0.33: Moderate, 0.67: Substantial (Chin, 1998)

0.25: Weak, 0.50: Moderate, 0.75: Substantial (Henseler et al., 2009)

< 0.30: Very weak, 0.30 - 0.50: Weak, 0.50 - 0.70: Moderate, 0.70 - 0.90: Strong, > 0.90: Very strong (Moore et al., 2013)

0.70: Generally considered good fit (Draper and Smith, 1998)

เกณฑ์ตัดสินใจสำหรับภาวะร่วมเส้นตรงพหุ (multicollinearity)

multicollinearity คือปัญหาที่ตัวแปรอิสระในสมการถดถอยมีความสัมพันธ์กันเองสูงมากเกินไป ถ้าตัวแปรอิสระมีความสัมพันธ์กันสูงจะมีผลให้ $V(\hat{\beta})$ มีค่าสูงมาก ผลคือค่า t-test สำหรับ $H: \beta \neq 0$ มีค่าต่ำมาก จึงทำให้ปฏิเสธตัวแปรอิสระเกือบทุกตัวหรือส่วนใหญ่ทั้งที่อาจเป็นตัวแปรสำคัญ

เกณฑ์ตัดสินใจสำหรับปัญหา multicollinearity มีหลายวิธี

1. สหสัมพันธ์ระหว่างตัวแปรสูงเกิน 0.85
2. sign, size, sig ของตัวแปรผิดปกติ
3. R_i^2 จาก Auxiliary Regression Model มีค่าสูงเกินกว่า 0.75
 - 1) $\text{Tolerance} = 1 - R_i^2 < 0.25$
 - 2) $\text{Variance Inflation Factor (VIF)} = \frac{1}{1 - R_i^2} > 4$

เกณฑ์ตัดสินใจสำหรับ beta

สปส. เส้นทางคือ สปส. การถดถอยที่เป็นน้ำหนักของตัวแปรสาเหตุ โดยปกติจะใช้เป็น standardized coefficient ถ้ามีค่าสูงแสดงว่าตัวแปรสาเหตุตัวนั้นมีอิทธิพลต่อตัวแปรผลลัพธ์สูง เกณฑ์ตัดสินใจมีหลายแนวทาง ดังนี้

$|\beta| < 0.10$: Small effect, $0.10 \leq |\beta| < 0.30$: Medium effect, $|\beta| \geq 0.30$: Large effect (Cohen, 1988).

$|\beta| \geq 0.20$: Meaningful effect (Chin, 1998).

$|\beta| < 0.10$: Small effect, $0.10 \leq |\beta| < 0.30$: Medium effect, $|\beta| \geq 0.30$: Large effect (Kline, 2011).

$|\beta| \geq 0.10$: Acceptable minimum. (Hair, Ringle, & Sarstedt, 2011).

$|\beta| > 0.50$: Strong effect (Henseler, Ringle, & Sinkovics, 2009).

$0.20 \leq |\beta| < 0.50$: Moderate effect, $|\beta| \geq 0.50$: Strong effect (Wetzels, Odekerken-Schröder, & Van Oppen, 2009).

$|\beta| \geq 0.10$: Relevant effect (Lohmöller, 1989).

เกณฑ์ตัดสินใจของ f-square

$f^2 < 0.02$: ปัจจัยสาเหตุมีอิทธิพลต่อปัจจัยผลลัพธ์ต่ำ

$0.02 \leq f^2 < 0.15$: ปัจจัยสาเหตุมีอิทธิพลต่อปัจจัยผลลัพธ์ปานกลาง

$0.15 \leq f^2 < 0.35$: ปัจจัยสาเหตุมีอิทธิพลต่อปัจจัยผลลัพธ์สูง

$f^2 \geq 0.35$: ปัจจัยสาเหตุมีอิทธิพลต่อปัจจัยผลลัพธ์สูงมาก (Cohen, 1988)

$$f^2 = \frac{R_{\text{included}}^2 - R_{\text{excluded}}^2}{1 - R_{\text{included}}^2}$$

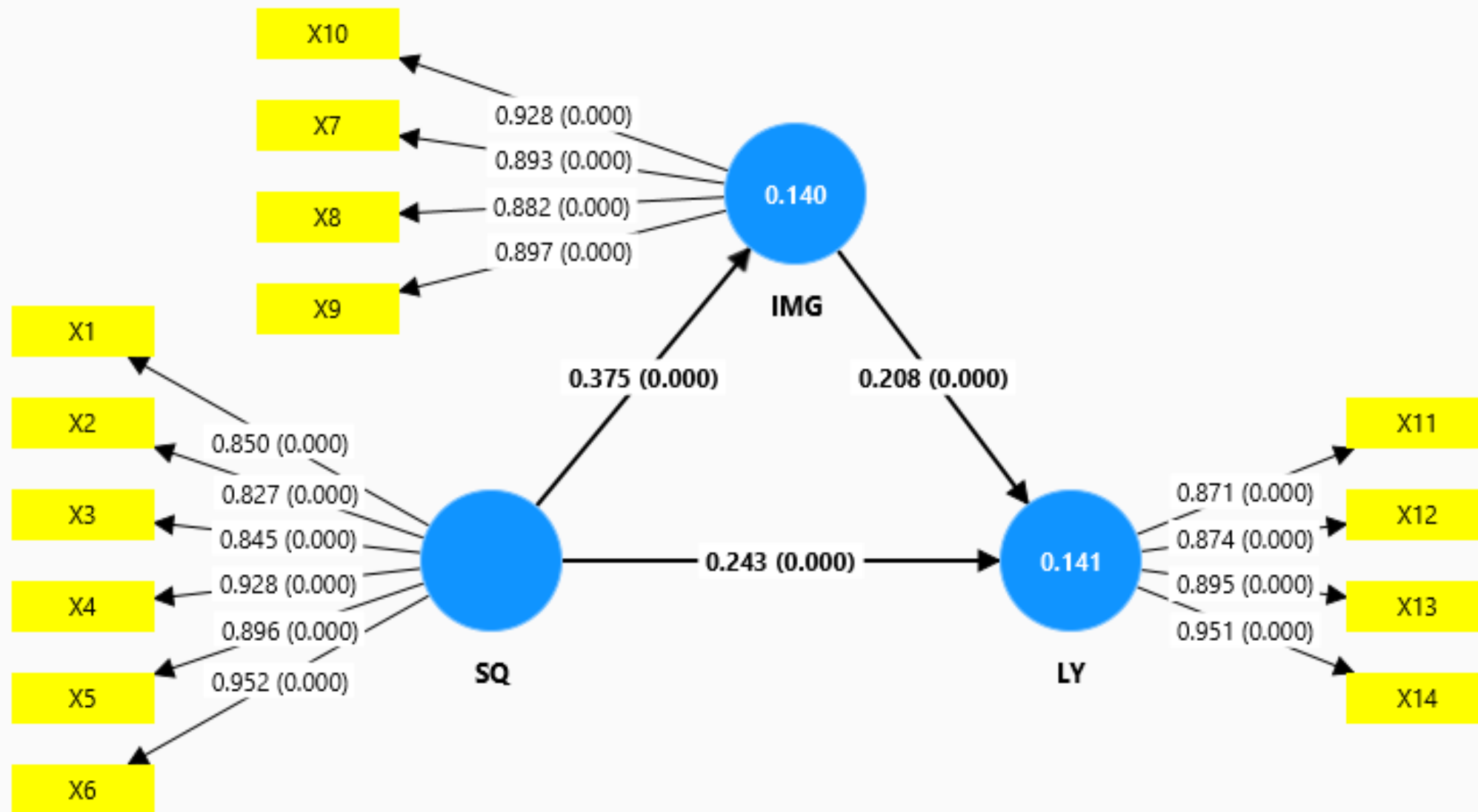
ให้เสนอค่า beta และ f^2 ในทุกเส้นทางซึ่งจะมีประโยชน์เพราะช่วยกันบอกเราว่าปัจจัยสาเหตุมีอิทธิพลระดับใด
ในกรณีตัวแปรตามที่ถูกกระทบจากตัวแปรอิสระหลายตัวจะสามารถอธิบายร่วมกับ สปส. เส้นทางได้ว่าตัวแปร
สาเหตุมีอิทธิพลลดหลั่นกันอย่างไร และแต่ละตัวมีอิทธิพลสูงต่ำเพียงใด

การวิเคราะห์ตัวแบบ

เรามีซอฟต์แวร์สำหรับวิเคราะห์ตัวแบบ SEM/Path model ให้เลือกใช้มากมาย วิธีเลือกใช้คือให้เลือกใช้ซอฟต์แวร์ที่มีความเป็นมิตรคือใช้งานง่าย

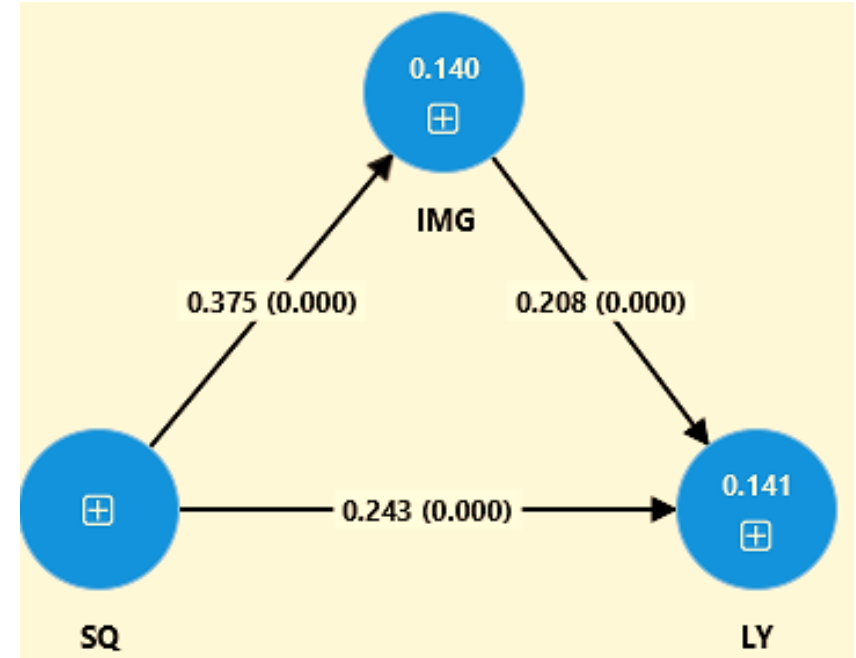
1. **กลุ่ม Variance-based SEM** กลุ่มนี้ใช้งานง่าย พัฒนาซอฟต์แวร์โดยใช้อัลกอริทึมที่ minimize variance โดยวิธี least square estimation เช่น Smart PLS, ADANCO, WARP และ PLS อื่นๆ ตัวอย่างต่อจากนี้จะรันด้วย Smart PLS 4
2. **กลุ่ม Covariance-based SEM** กลุ่มนี้พัฒนาซอฟต์แวร์โดยใช้อัลกอริทึม maximum likelihood estimation โดย maximize เทียบกับ covariance matrix เช่น AMOS, LISREL
3. ซอฟต์แวร์ทางสถิติ เช่น R, Strata
4. **PROCESS macro** ซอฟต์แวร์นี้มุ่งเน้นที่การวิเคราะห์ mediation model, moderation model, moderated mediation model ที่มีตัวแปรสาเหตุตัวเดียว มีตัวแปรผลลัพธ์ตัวเดียว ก่อนใช้ต้องแปลงข้อมูล multi-valued variable/composite variable เป็น single-valued variable เสียก่อน การวิเคราะห์จะเสนอเป็น unstandardized coefficient แต่ก็แสดงค่า standardized coefficient ถ้าต้องการโดยขอผ่านไคอะล็อก

ตัวอย่าง Causation model



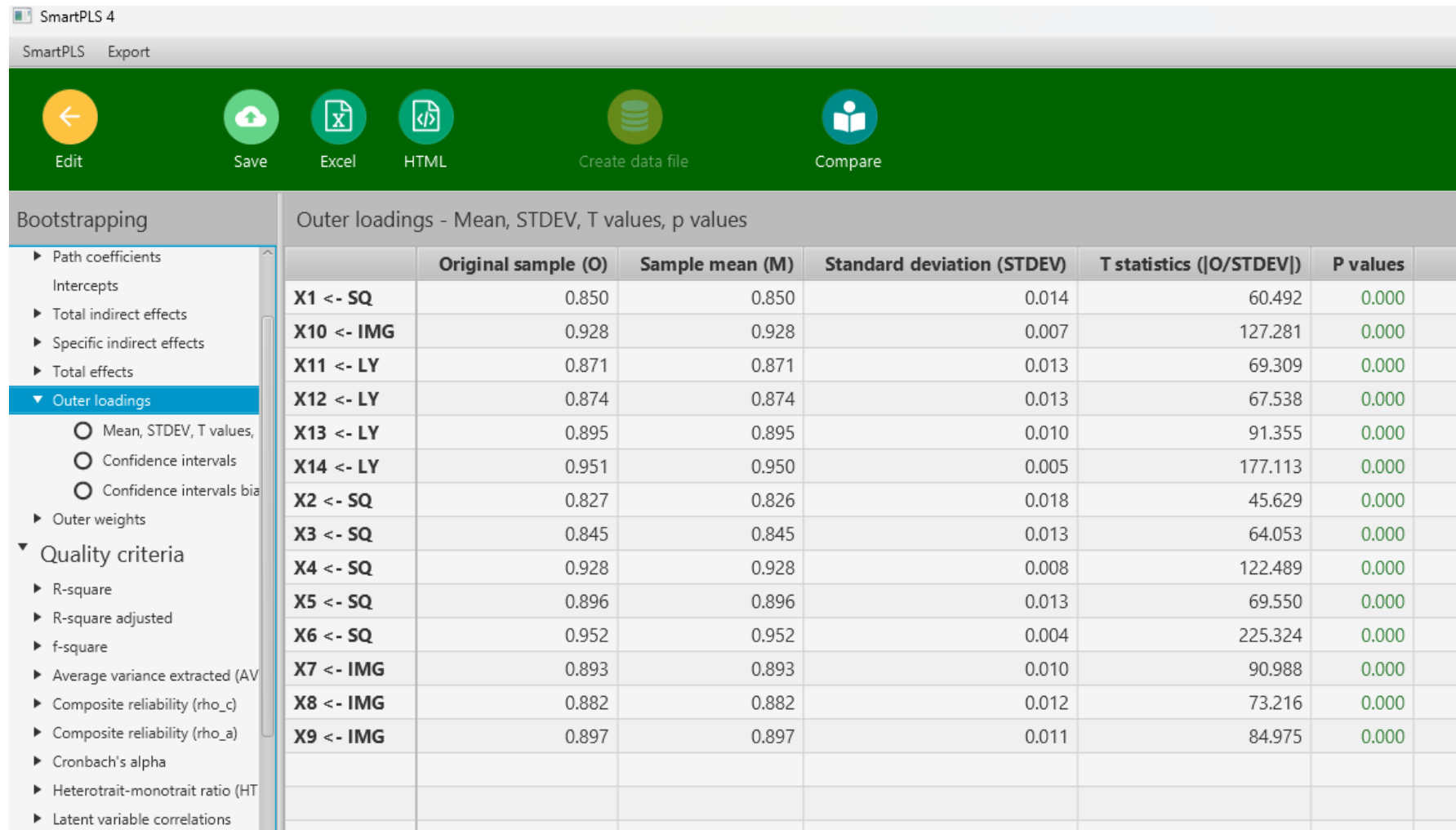
สรุปผลการวิเคราะห์

1. ผลการวิเคราะห์คุณภาพเครื่องมือพบว่า loading มีค่าสูงกว่า 0.707 เป็นปริมาณบวกและมีนัยสำคัญทุกค่า แสดงว่ามีความตรงเชิงเหมือน HTMT ของตัวแปรแฝงทุกคู่มีค่าต่ำกว่า 0.85 และมีนัยสำคัญ ($H: HTMT < .85$) แสดงว่ามาตรวัดมีความเที่ยงตรงเชิงจำแนก มาตรวัดมีค่า $CR > 0.60$, Cronbach's Alpha > 0.70 , AVE > 0.5 ทุกตัวแปร แสดงว่ามาตรวัดมีความเที่ยง
2. ผลการวิเคราะห์ตัวแบบพบว่าอิทธิพลตามเส้นทางทุกเส้นทางมีค่าสูง (สูงกว่า 0.30) เป็นปริมาณบวกตรงตามบริบท และมีนัยสำคัญ แสดงว่าทุกสมมติฐานเป็นจริง
3. ผลการวิเคราะห์แสดงให้เห็นว่าข้อมูลที่ได้รับจากการสำรวจในพื้นที่เป้าหมายสนับสนุนความจริงทางทฤษฎีว่าปัจจัยสาเหตุมีอิทธิพลต่อปัจจัยผลลัพธ์จริงคือ $SQ \rightarrow LY$, $SQ \rightarrow IMF$ และ $IMG \rightarrow LY$ ในบริบทที่เปลี่ยนไปจากเดิมของทฤษฎีเหล่านั้น
4. เมื่อพิจารณาอิทธิพลของตัวแปร SQ และ IMG ที่พบว่ามอิทธิพลต่อ LY นั้นจะเห็นได้ว่า SQ มีอิทธิพลสูงกว่า IMG ประกอบกับการลงทุนเพื่อสร้าง SQ ถูกกว่าการสร้าง IMG มากและใช้เวลาน้อยกว่ามาก ดังนั้นผู้ประกอบการต้องเร่งรัดพนักงานให้มุ่งที่ SQ อย่างเต็มกำลัง ผลที่ตามมาคือความภักดีที่แข็งแกร่งยั่งยืนของลูกค้า ขณะเดียวกันก็ให้ดำเนินการเพื่อสร้างภาพลักษณ์ที่ดีผ่านสื่อและ WOM



ผลการวิเคราะห์ข้อมูลจากการสำรวจ

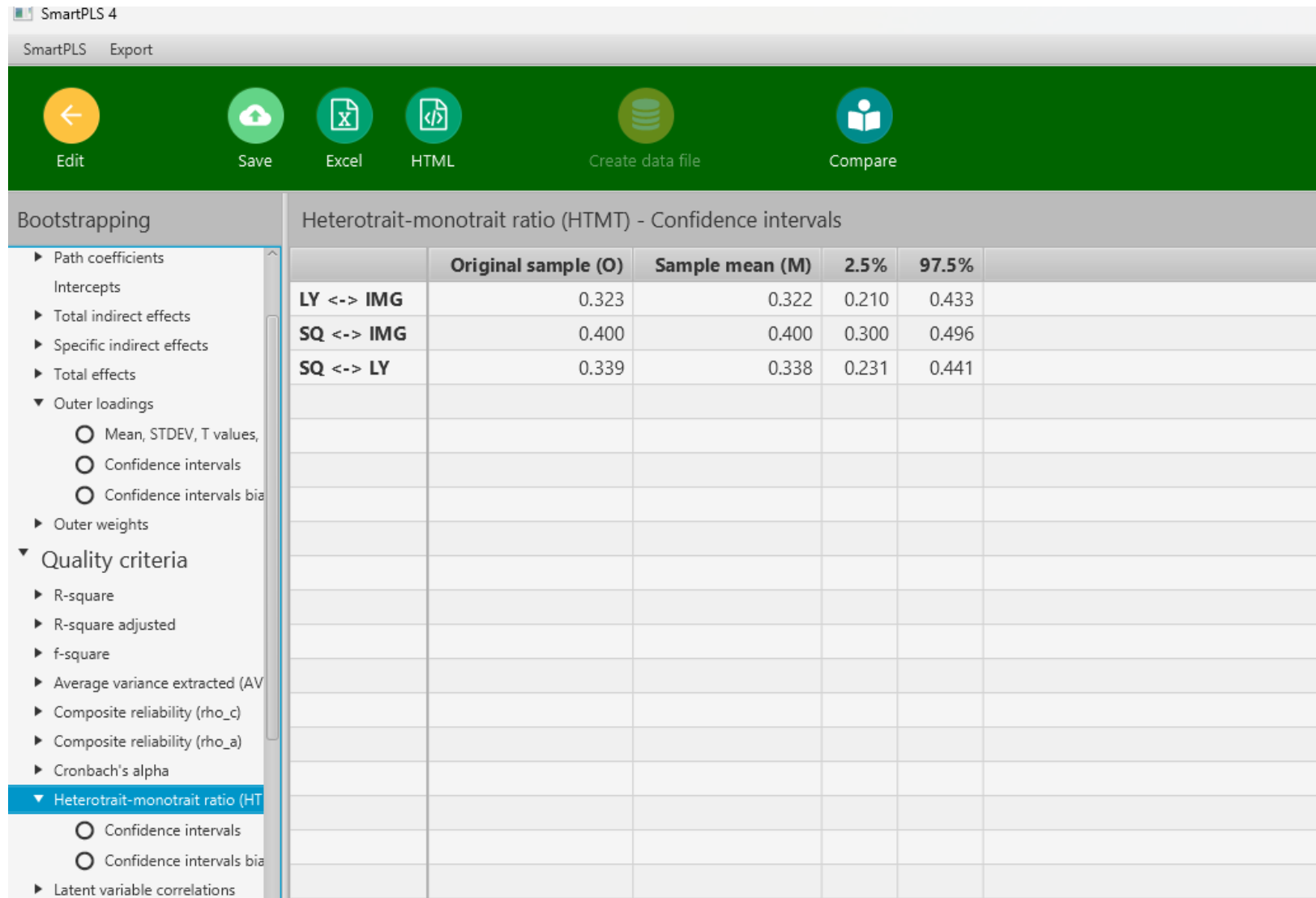
1. Convergent validity



	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
X1 <- SQ	0.850	0.850	0.014	60.492	0.000
X10 <- IMG	0.928	0.928	0.007	127.281	0.000
X11 <- LY	0.871	0.871	0.013	69.309	0.000
X12 <- LY	0.874	0.874	0.013	67.538	0.000
X13 <- LY	0.895	0.895	0.010	91.355	0.000
X14 <- LY	0.951	0.950	0.005	177.113	0.000
X2 <- SQ	0.827	0.826	0.018	45.629	0.000
X3 <- SQ	0.845	0.845	0.013	64.053	0.000
X4 <- SQ	0.928	0.928	0.008	122.489	0.000
X5 <- SQ	0.896	0.896	0.013	69.550	0.000
X6 <- SQ	0.952	0.952	0.004	225.324	0.000
X7 <- IMG	0.893	0.893	0.010	90.988	0.000
X8 <- IMG	0.882	0.882	0.012	73.216	0.000
X9 <- IMG	0.897	0.897	0.011	84.975	0.000

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

2. Discriminant validity



The screenshot shows the SmartPLS 4 software interface. The top navigation bar includes icons for Edit, Save, Excel, HTML, Create data file, and Compare. The left sidebar lists various analysis options under 'Bootstrapping' and 'Quality criteria'. The 'Heterotrait-monotrait ratio (HTMT) - Confidence intervals' table is displayed on the right, showing results for three pairs of variables: LY <-> IMG, SQ <-> IMG, and SQ <-> LY. The table includes columns for the original sample (O), sample mean (M), and 2.5% and 97.5% confidence intervals.

	Original sample (O)	Sample mean (M)	2.5%	97.5%
LY <-> IMG	0.323	0.322	0.210	0.433
SQ <-> IMG	0.400	0.400	0.300	0.496
SQ <-> LY	0.339	0.338	0.231	0.441

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

3. Reliability

Average variance extracted (AVE) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG	0.811	0.810	0.011	72.130	0.000	
LY	0.807	0.807	0.011	70.471	0.000	
SQ	0.782	0.782	0.013	61.735	0.000	

Composite reliability (rho_c) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG	0.945	0.945	0.004	246.527	0.000	
LY	0.944	0.943	0.004	239.535	0.000	
SQ	0.955	0.955	0.003	299.156	0.000	

Composite reliability (rho_a) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG	0.923	0.926	0.006	149.027	0.000	
LY	0.927	0.930	0.007	132.595	0.000	
SQ	0.949	0.951	0.005	199.947	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Cronbach's alpha - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG	0.922	0.922	0.006	161.049	0.000	
LY	0.920	0.920	0.006	155.510	0.000	
SQ	0.944	0.944	0.004	225.126	0.000	

Total effects - Mean, STDEV, T values, p values

Copy to Excel/Word

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG -> LY	0.208	0.209	0.052	3.982	0.000	
SQ -> IMG	0.375	0.377	0.046	8.080	0.000	
SQ -> LY	0.321	0.322	0.050	6.434	0.000	

f-square - Mean, STDEV, T values, p values

Copy to Excel/Word

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG -> LY	0.043	0.047	0.023	1.847	0.065	
SQ -> IMG	0.163	0.171	0.049	3.365	0.001	
SQ -> LY	0.059	0.062	0.027	2.153	0.031	

ตัวอย่าง Parallel mediation analysis

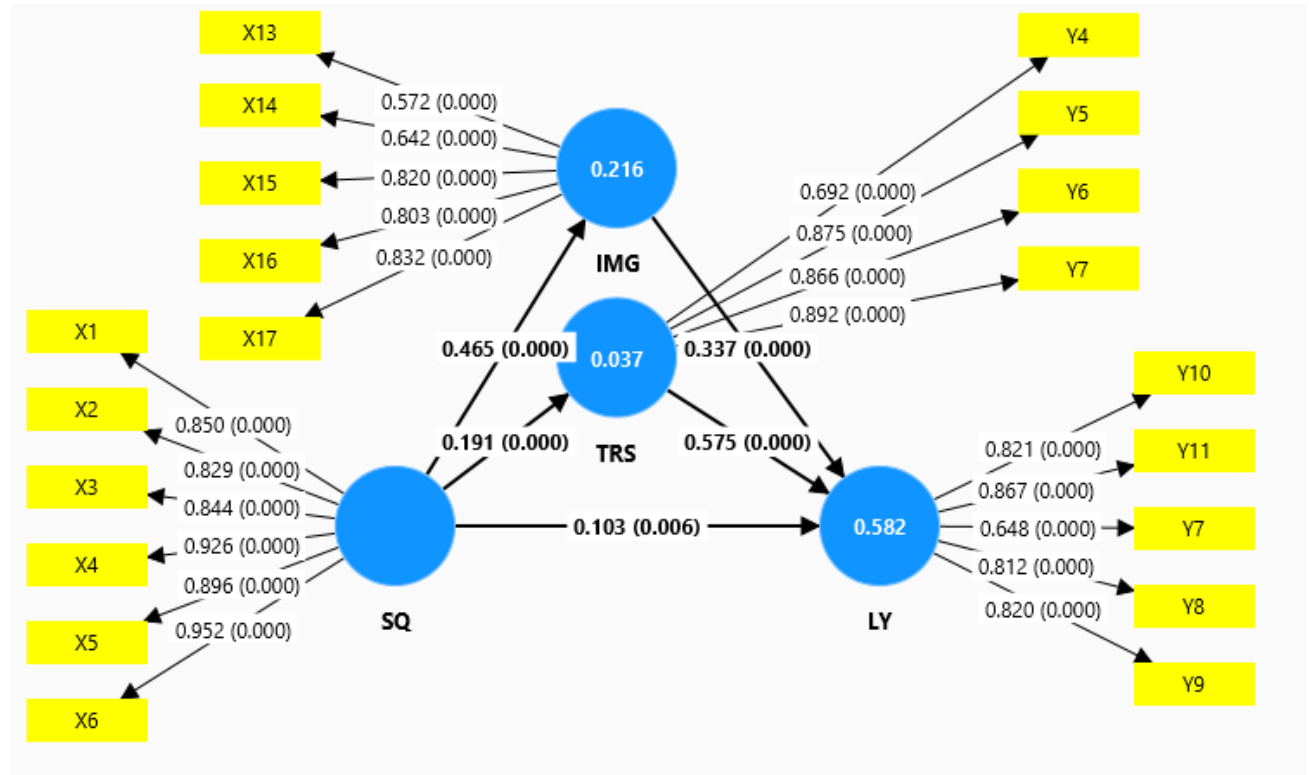
Mediation model มีจุดประสงค์เพื่อศึกษาว่าปัจจัยสาเหตุส่งผลกระทบต่อปัจจัยผลลัพธ์อย่างไร (HOW) โดยตรวจสอบว่าอิทธิพลรวม (Total Effect, TE) คืออิทธิพลตามเส้นทาง $X \rightarrow Y$ ที่มีค่าสูงมากนั้นเป็นเพราะอิทธิพลของปัจจัยสาเหตุหรือยังมีปัจจัยอื่นแฝงตัวเชื่อมโยงเอาไว้

วิธีการศึกษา mediation analysis จะเริ่มต้นจากการหาค่าตัวแปรนั้นจากผลการสัมภาษณ์เชิงลึก หรือจากข้อเสนอแนะจากงานวิจัยอื่น หรือมีทฤษฎีใดชี้ทางเอาไว้ หรือจากความสงสัยว่าอิทธิพลรวม (TE) สูงเกินไปคือ $TE > 0.30, 0.50$ (Cohen, 1988)

ตัวแปรต้นกลางอาจมีเพียงตัวเดียวหรือหลายตัว ถ้ามีหลายตัวอาจเป็นแบบขนาน parallel mediation) ดังภาพขวามือ หรือแบบอนุกรม (serial mediation)

การวิเคราะห์เราจะตรวจสอบว่าอิทธิพลรวมเดิมนั้นสำคัญหรือไม่ และเมื่อได้แทรกตัวแปรต้นกลางไปแล้วอิทธิพล $X \rightarrow Y$ เรียกชื่อใหม่ว่า Direct Effect (DE) ลดลงมากเพียงใดยังมีความสำคัญอยู่หรือไม่ อิทธิพลทางอ้อม (indirect effect, IE) แต่ละเส้นทางจากตัวแปรสาเหตุผ่านตัวแปรต้นกลางแล้วต่อไปยังตัวแปรผลลัพธ์มีค่าเท่าไร มีนัยสำคัญหรือไม่ เส้นทางใดมีความสำคัญกว่า

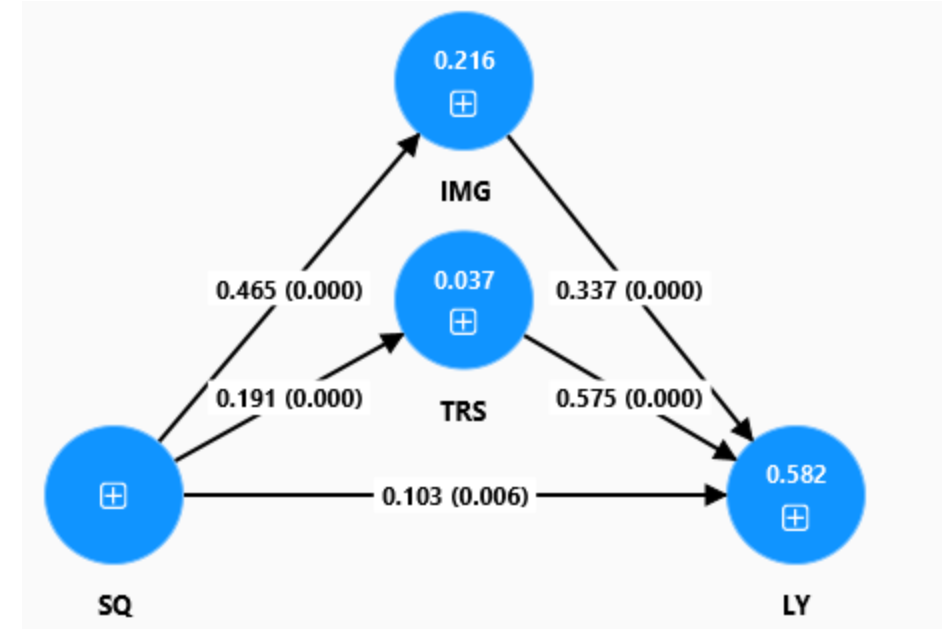
จากภาพทางขวามือ $TE = 0.370^{**}$, $DE = 0.103^{**}$ เส้นทาง
 $SQ \rightarrow IMG \rightarrow LY$ มี $IE1 = 0.465 * 0.337 = 0.157^{**}$ เส้นทาง
 $SQ \rightarrow TRS \rightarrow LY$ มี $IE2 = 0.191 * 0.575 = 0.110^{**}$
 $TIE = 0.267^{**}$



สรุปผลการวิเคราะห์ข้อมูล

ผลการวิเคราะห์สรุปได้ดังนี้

1. ผลการประเมินคุณภาพเครื่องมือพบว่ามาตรวัดมี construct validity คือมีความตรงเชิงเหมือน ตัวชี้วัดมีค่า loadings > 0.707 มีอยู่ 3 ค่าที่ต่ำกว่านี้แต่ก็สูงกว่า 0.50 และมี AVE > 0.50 มีค่าเป็นบวกและมีนัยสำคัญ ตัวชี้วัดไม้ไขว้กลุ่มสังเกตจากค่า HTMT อยู่ระหว่าง 0.218 ถึง 0.709 ไม่เกิน 0.85 มีความเที่ยงอยู่ในเกณฑ์คือ AVE > 0.50, Cronbach's alpha > 0.70, Rho_a > 0.60, Rho_c > 0.60
2. TE = 0.370**, DE = 0.103** แสดงว่าที่จริงแล้ว SQ มีอิทธิพลต่อ LY ไม่มากคือมีอิทธิพลเพียง 28% อิทธิพลที่มากถึง 72% ก็คืออิทธิพลของตัวแปรคนกลาง พบว่าอิทธิพลตามเส้นทาง SQ → IMG → LY มีค่าเท่ากับ 0.157* และอิทธิพลตามเส้นทาง SQ → TRS → LY มีค่าเท่ากับ 0.110* อิทธิพลทางอ้อมรวมเท่ากับ 0.267** แสดงว่า TRS และ IMG มีอิทธิพลในฐานะปัจจัยเชื่อมโยงแต่ IMG มีความสำคัญสูงกว่า
3. ในชั้นปฏิบัติให้ผู้ประกอบการใส่ใจคุณภาพบริการตามปกติอย่าได้ปล่อยปละ แต่ปัจจัยที่ต้องให้ความสำคัญมากกว่าคือภาพลักษณ์และความพึงพอใจของผู้บริโภค



ตัวอย่างมาตรวัด

Tourist Satisfaction scale

โดยรวมแล้ว ฉันพอใจกับประสบการณ์การท่องเที่ยวครั้งนี้
การบริการท่องเที่ยวตรงกับความคาดหวังของฉัน
ฉันรู้สึกคุ้มค่ากับเงินที่จ่ายไป
กิจกรรมและสถานที่ท่องเที่ยวที่น่าประทับใจ
มัคคุเทศก์และเจ้าหน้าที่ให้บริการอย่างเป็นมิตร
อาหารและที่พักมีคุณภาพดี
สถานที่ท่องเที่ยวมีความปลอดภัย
การเดินทางสะดวกและราบรื่น
ฉันได้เรียนรู้วัฒนธรรมใหม่ๆ จากการเดินทางครั้งนี้
มีสิ่งอำนวยความสะดวกครบครัน
เวลาในการให้บริการเหมาะสม
พนักงานมีความเชี่ยวชาญและน่าเชื่อถือ
ฉันรู้สึกผ่อนคลายและมีความสุขระหว่างการเดินทาง
ประสบการณ์ในครั้งนี้เกินความคาดหมาย
ฉันอยากกลับมาเที่ยวที่นี่อีกครั้ง

Destination Image Scale

สถานที่ท่องเที่ยวมีความสวยงามตามธรรมชาติ
มีความหลากหลายทางวัฒนธรรม
สถานที่สะอาดและเป็นระเบียบ
มีความปลอดภัยสำหรับนักท่องเที่ยว
สภาพภูมิอากาศเหมาะแก่การท่องเที่ยว
มีเอกลักษณ์เฉพาะตัวที่น่าดึงดูด
สามารถเดินทางได้สะดวก
มีสถานที่ท่องเที่ยวที่น่าสนใจหลากหลาย
ประชาชนในพื้นที่เป็นมิตร
มีสิ่งอำนวยความสะดวกครบครัน
มีอาหารและของกินที่น่าสนใจ
มีราคาที่เหมาะสม
มีความน่าสนใจทางประวัติศาสตร์

ผลการวิเคราะห์ข้อมูลจากการสำรวจ-Total Indirect Effect

Total indirect effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
SQ -> LY	0.267	0.268	0.037	7.275	0.000	

Specific indirect effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
SQ -> IMG -> LY	0.157	0.157	0.024	6.550	0.000	
SQ -> TRS -> LY	0.110	0.110	0.028	3.867	0.000	

Total effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG -> LY	0.337	0.338	0.041	8.304	0.000	
SQ -> IMG	0.465	0.466	0.043	10.891	0.000	
SQ -> LY	0.370	0.371	0.046	7.997	0.000	
SQ -> TRS	0.191	0.193	0.050	3.828	0.000	
TRS -> LY	0.575	0.575	0.035	16.202	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ-Outer loadings

Outer loadings - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
X1 <- SQ	0.850	0.850	0.014	61.846	0.000	
X13 <- IMG	0.572	0.570	0.056	10.206	0.000	
X14 <- IMG	0.642	0.640	0.050	12.737	0.000	
X15 <- IMG	0.820	0.819	0.024	34.679	0.000	
X16 <- IMG	0.803	0.802	0.026	30.323	0.000	
X17 <- IMG	0.832	0.832	0.021	39.777	0.000	
X2 <- SQ	0.829	0.829	0.017	48.427	0.000	
X3 <- SQ	0.844	0.844	0.013	65.296	0.000	
X4 <- SQ	0.926	0.926	0.008	120.990	0.000	
X5 <- SQ	0.896	0.895	0.012	73.877	0.000	
X6 <- SQ	0.952	0.952	0.004	233.433	0.000	
Y10 <- LY	0.821	0.821	0.019	42.442	0.000	
Y11 <- LY	0.867	0.866	0.015	56.076	0.000	
Y4 <- TRS	0.692	0.692	0.033	21.253	0.000	
Y5 <- TRS	0.875	0.874	0.013	66.687	0.000	
Y6 <- TRS	0.866	0.866	0.014	63.260	0.000	
Y7 <- LY	0.648	0.648	0.032	20.417	0.000	
Y7 <- TRS	0.892	0.892	0.010	89.500	0.000	
Y8 <- LY	0.812	0.812	0.021	38.781	0.000	
Y9 <- LY	0.820	0.820	0.019	42.148	0.000	

Loading > .707,
positive and significant
(H0: $\lambda = 0$ vs H1: $\lambda > 0$)

ผลการวิเคราะห์ข้อมูลจากการสำรวจ-Heterotrait-Monotrait (HTMT)

Heterotrait-monotrait ratio (HTMT) - Confidence intervals					
	Original sample (O)	Sample mean (M)	2.5%	97.5%	
LY <-> IMG	0.611	0.611	0.527	0.693	
SQ <-> IMG	0.533	0.531	0.432	0.623	
SQ <-> LY	0.419	0.418	0.317	0.512	
TRS <-> IMG	0.224	0.229	0.122	0.342	
TRS <-> LY	0.709	0.710	0.637	0.779	
TRS <-> SQ	0.218	0.218	0.116	0.325	

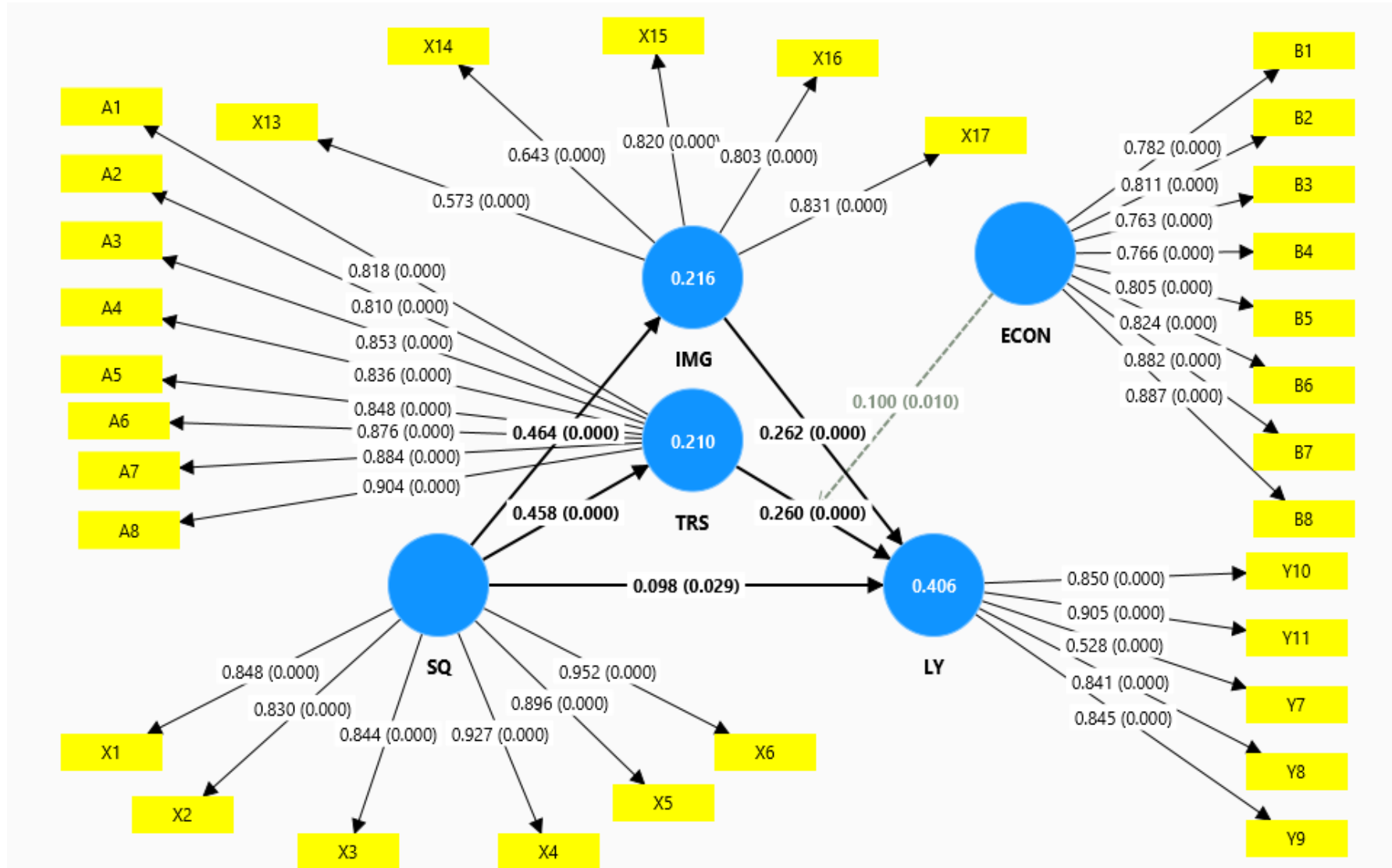
Ho: HTMT $\geq$ 0.85/0.90 (lack of discriminant validity) vs H1: HTMT < 0.85/0.90
(discriminant validity exists)

ผลการวิเคราะห์ข้อมูลจากการสำรวจ ค่า f-square

f-square - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
IMG -> LY	0.211	0.217	0.056	3.767	0.000	
SQ -> IMG	0.275	0.284	0.066	4.162	0.000	
SQ -> LY	0.020	0.022	0.015	1.343	0.179	
SQ -> TRS	0.038	0.042	0.021	1.774	0.076	
TRS -> LY	0.753	0.765	0.120	6.298	0.000	

Benchmarks คือ 0.02, 0.15, 0.35 (Cohen, 1988)

ตัวอย่าง Moderated mediation model



ตัวอย่าง Moderated mediation model

1. ผลการวิเคราะห์คุณภาพเครื่องมือ พบว่า loading มีค่าตรงเกณฑ์คือไม่ต่ำกว่า 0.707 มีค่าบวก และมีนัยสำคัญ (ยกเว้น X14 และ Y7 แต่ก็สูงกว่า 0.5 และสอดคล้องเงื่อนไขที่ Hair (2018) แนะนำ HTMT อยู่ในช่วง 0.347-0.515, $CR > 0.60$ $\alpha > 0.70$, $AVE > 0.50$ แสดงว่ามาตรวัดมีคุณภาพ

2. ผลการวิเคราะห์อิทธิพลทางตรงและอิทธิพลทางอ้อม $TE = 0.338^{***}$, $DE = 0.098^*$, $IE1: SQ \rightarrow TRS \rightarrow LY = 0.119^{**}$ $IE2: SQ \rightarrow IMG \rightarrow LY = 0.122^{***}$

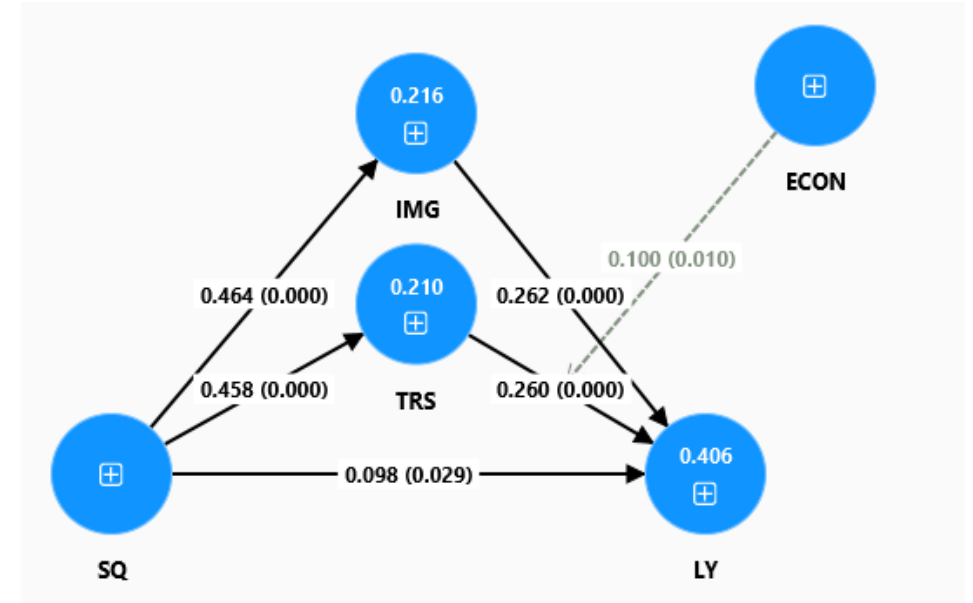
แสดงว่า SQ มีได้มีอิทธิพลต่อ LY สูงมากตามค่า TE ($\beta = 0.338$) ที่ปรากฏค่าสูงมากเช่นนี้เป็นอิทธิพลหลอกเพราะมี TRS และ IMG ซ่อนตัวเชื่อมโยงอยู่ สังเกตได้ว่าเมื่อแทรก TRS และ IMG คั่นกลางระหว่าง SQ และ LY พบว่าอิทธิพลทางตรงเหลือเพียง 0.098 ($p = 0.029$) ลดลงถึงร้อยละ 71 หรือ SQ มีอิทธิพลต่อ LY ร้อยละ 29 เมื่อพิจารณาอิทธิพลทางอ้อมพบว่า IE1 และ IE2 มีอิทธิพลใกล้เคียงกัน

3. ผลการวิเคราะห์ IE1: $SQ \rightarrow TRS \rightarrow LY$ ตามเงื่อนไขของระดับค่าของ ECON พบว่า ECON มีนัยสำคัญ ($\beta = 0.100$, $p = 0.010$) พบอิทธิพลทางอ้อม IE1: $SQ \rightarrow TRS \rightarrow LY$ มีค่าสูงขึ้นตาม ECON ที่มีค่าสูงขึ้นคือ

CIE1: IE1: $SQ \rightarrow TRS \rightarrow LY = 0.061^*$ เมื่อ ECON มีค่าต่ำ

CIE1: IE1: $SQ \rightarrow TRS \rightarrow LY = .102^{***}$ เมื่อ ECON มีค่าปานกลาง

CIE1: IE1: $SQ \rightarrow TRS \rightarrow LY = 0.143^{**}$ เมื่อ ECON มีค่าสูง



สรุปผลการวิเคราะห์ Moderated mediation model

4. สรุปผล (conclusion and implication)

ผลการวิเคราะห์ตัวแบบสรุปได้ว่า SQ มีอิทธิพลทางตรงต่อ LY แต่มีอิทธิพลต่ำมากแม้ว่าจะมีนัยสำคัญ ($\beta = 0.098$, $p = 0.029$) (Cohen, 1988; Kline, 2011) คิดเป็นเพียงร้อยละ 29 อิทธิพลหลักคืออิทธิพลทางอ้อม

อิทธิพลทางอ้อมที่ผ่านมาทาง IMG สูงกว่าผ่านทาง TRS เล็กน้อย TRS (Tourist Satisfaction) เกี่ยวข้องกับที่พัก สิ่งอำนวยความสะดวก บริการเที่ยวบิน/รถทัวร์ การขนส่งในท้องถิ่น บริการอาหารและเครื่องดื่ม สถานที่ท่องเที่ยวและบริการไกด์นำเที่ยว บริการช้อปปิ้ง คุณภาพสินค้าและบริการและความยุติธรรมของราคา ความคุ้มค่าเงิน/เวลา และการสนองความคาดหวัง ขณะที่ IMG (Destination Image) มีความเกี่ยวข้องกับคุณภาพโครงสร้างพื้นฐาน ความสวยงามของธรรมชาติ มรดกทางวัฒนธรรม ระดับราคา สภาพอากาศ ความเป็นเอกลักษณ์เมื่อเทียบกับที่อื่น และความรู้สึกทางอารมณ์ต่อจุดหมายปลายทาง IMG จึงทำให้ดีขึ้นได้ยากกว่า TRS ดังนั้นพื้นที่ปลายทางควรเน้นไปที่ TRS ก่อน ทำให้มากเท่าที่พึงเป็นไปได้ ขณะเดียวกันก็ค่อยๆ พัฒนาภาพลักษณ์พื้นที่ปลายทางไปเรื่อย ๆ เช่นอนุรักษ์ป่าเขา แหล่งน้ำ ทะเล อนุรักษ์ประเพณีท้องถิ่น สร้างความสำนึกเรื่องความเป็นมิตรกับผู้มาเยือนอย่างต่อเนื่อง

สำหรับ TRS ยังพบว่า ถ้า ECON สูงขึ้นๆ อิทธิพลของ SQ ที่มีต่อ LY จะสูงขึ้นตาม แสดงว่าพื้นที่ปลายทางต้องเน้นการสร้างความสะดวกสบาย ความคุ้มค่า ให้แก่นักท่องเที่ยวในประสบการณ์เชิงสุนทรีย์ (Esthetic Experience) ประสบการณ์การหนีหนี (Escapist Experience) ประสบการณ์การเรียนรู้ (Educational Experience) และประสบการณ์ความบันเทิง (Entertainment Experience) ซึ่งจะมีผลให้ความภักดีต่อพื้นที่ปลายทางสูงขึ้นจนมีการกล่าวถึงในทางที่ดี บอกต่อ และมาท่องเที่ยวซ้ำๆ

Experience Economy

หมายเหตุ Experience Economy คือเศรษฐกิจที่สร้างคุณค่าผ่านประสบการณ์และความจดจำของลูกค้าแทนที่จะขายแต่สินค้าหรือบริการ โดยนำ big data, AI, Cloud และ IOT มาช่วยวิเคราะห์เพื่อสามารถอ่านใจลูกค้าได้ โดยส่งข้อมูลจากบิ๊กดาต้าให้ AI วิเคราะห์แล้วส่งผลการวิเคราะห์ไปให้ IOT ทำงานแบบอัตโนมัติแล้วส่งไปประเมินผลบน cloud โดยเอไอจะวิเคราะห์เจาะจงเฉพาะตัวลูกค้า เน้นความพึงพอใจซึ่งจะส่งผลต่อเนื่องถึงความภักดี

Experience economy มีสมบัติคือ memorability (ประทับใจ ผั่งใจ ตรึงใจ มีช่วงเวลาที่น่าจดจำ เป็นประสบการณ์ที่ทำให้เกิดความทรงจำที่แรงกล้า ติดตามไปนาน และมีความหมายพิเศษ) Immersive (ดื่มด่ำ จดจ่อกับการมีส่วนร่วมอย่างลึกซึ้ง เช่นร่วมออกแบบผลิตภัณฑ์ ร่วมทำอาหารกับเชฟ จดจ่อกับประสบการณ์ เป็นการมีส่วนร่วมอย่างลึกซึ้ง จดจ่อเต็มที่ ไม่ใช่แค่ชมหรือฟังเฉยๆ) Premium price (ประสบการณ์บริโภคผลิตภัณฑ์/บริการในราคาสูง Premium price ไม่ใช่แค่การขายแพง แต่เป็น "การรับประกันว่าจะได้ประสบการณ์ที่ไม่ผิดหวัง, ไม่ซ้ำใคร, และมีคุณภาพสูงสุด) Sharable moment (ส่งต่อความประทับใจ ความดื่มด่ำ และความจดจำผ่านสื่อสังคมออนไลน์) Personal transformation (เปลี่ยนแปลงทักษะ เปลี่ยน mindset) language immersion (ประสบการณ์ที่ทำให้ได้เรียนรู้และใช้ภาษาใหม่ผ่านกิจกรรมจริง) wellness retreat (ประสบการณ์ที่มุ่งเน้นการฟื้นฟูสุขภาพจิตใจ ลดความเครียด เพิ่มความมั่นใจผ่านกิจกรรมการฝึกอบรมจิต)

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Total effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON -> LY	0.200	0.201	0.052	3.861	0.000	
IMG -> LY	0.262	0.263	0.050	5.208	0.000	
SQ -> IMG	0.464	0.465	0.043	10.850	0.000	
SQ -> LY	0.338	0.338	0.048	6.988	0.000	
SQ -> TRS	0.458	0.459	0.044	10.351	0.000	
TRS -> LY	0.260	0.260	0.051	5.118	0.000	
ECON x TRS -> LY	0.100	0.099	0.039	2.576	0.010	

Specific indirect effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
SQ -> TRS -> LY	0.119	0.119	0.027	4.481	0.000	
SQ -> IMG -> LY	0.122	0.122	0.027	4.578	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Conditional indirect effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
SQ -> TRS -> LY conditional on ECON at +1 SD	0.143	0.142	0.026	5.588	0.000	
SQ -> TRS -> LY conditional on ECON at -1 SD	0.061	0.061	0.028	2.151	0.032	
SQ -> TRS -> LY conditional on ECON at Mean	0.102	0.102	0.022	4.636	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Average variance extracted (AVE) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.666	0.666	0.018	37.341	0.000	
IMG	0.550	0.550	0.020	27.562	0.000	
LY	0.648	0.649	0.016	39.807	0.000	
SQ	0.782	0.782	0.013	62.118	0.000	
TRS	0.730	0.729	0.016	45.475	0.000	

Composite reliability (rho_c) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.941	0.941	0.005	208.815	0.000	
IMG	0.857	0.856	0.011	78.584	0.000	
LY	0.900	0.900	0.007	129.658	0.000	
SQ	0.955	0.955	0.003	301.256	0.000	
TRS	0.956	0.956	0.003	275.573	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Composite reliability (rho_a) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.933	0.935	0.006	168.712	0.000	
IMG	0.809	0.812	0.019	43.447	0.000	
LY	0.879	0.880	0.011	77.259	0.000	
SQ	0.949	0.950	0.004	227.544	0.000	
TRS	0.950	0.951	0.004	233.499	0.000	
ECON x TRS	1.000	1.000	0.000	n/a	n/a	

Cronbach's alpha - Mean, STDEV, T values, p values

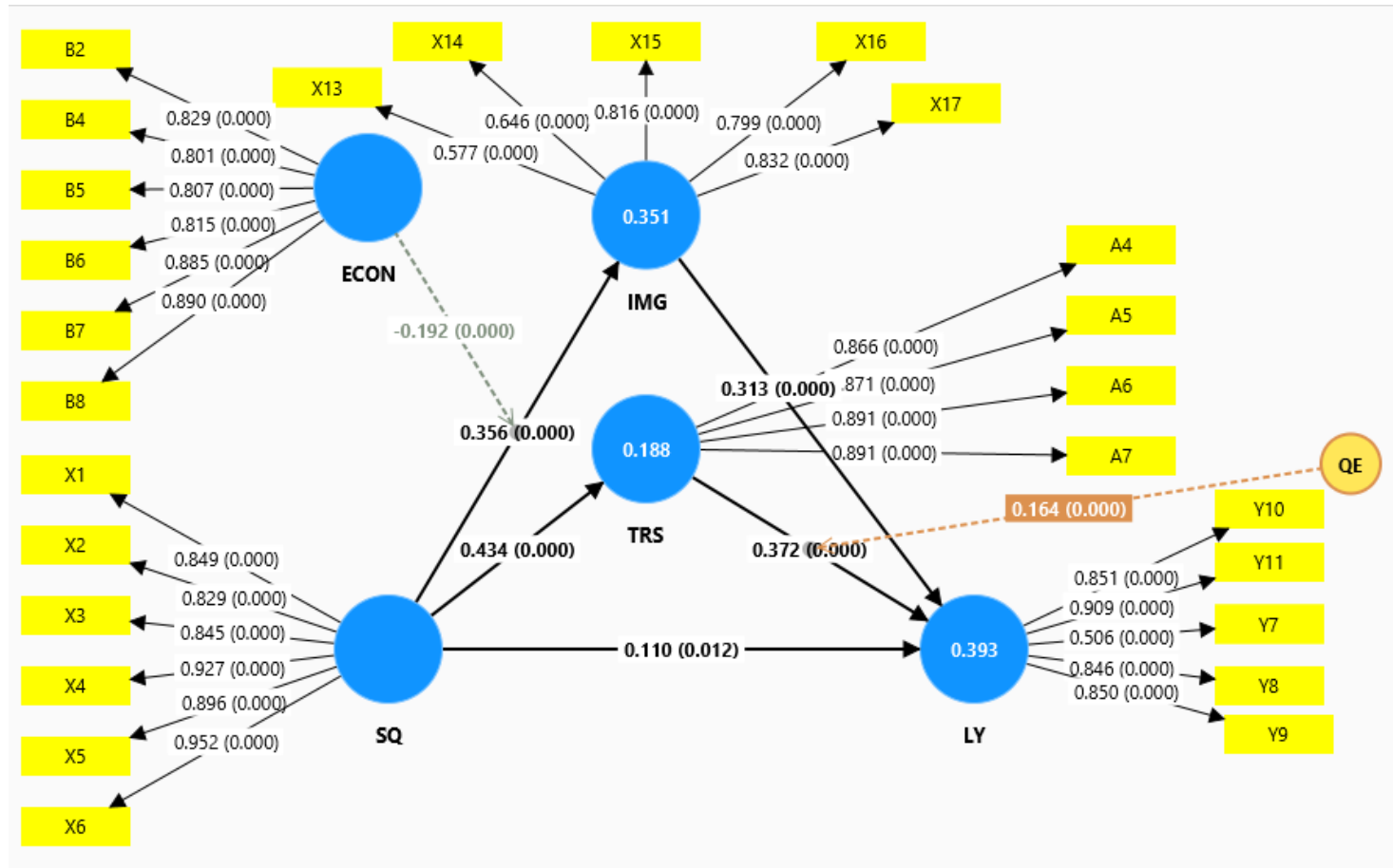
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.928	0.928	0.006	159.517	0.000	
IMG	0.790	0.790	0.018	42.876	0.000	
LY	0.855	0.855	0.012	73.970	0.000	
SQ	0.944	0.944	0.004	225.126	0.000	
TRS	0.947	0.947	0.004	218.230	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Heterotrait-monotrait ratio (HTMT) - Confidence intervals					
	Original sample (O)	Sample mean (M)	2.5%	97.5%	
IMG <-> ECON	0.515	0.514	0.412	0.610	
LY <-> ECON	0.534	0.533	0.439	0.621	
LY <-> IMG	0.611	0.611	0.527	0.693	
SQ <-> ECON	0.347	0.346	0.240	0.445	
SQ <-> IMG	0.533	0.531	0.432	0.623	
SQ <-> LY	0.419	0.418	0.317	0.512	
TRS <-> ECON	0.520	0.519	0.431	0.603	
TRS <-> IMG	0.595	0.594	0.509	0.671	
TRS <-> LY	0.592	0.592	0.517	0.663	
TRS <-> SQ	0.480	0.479	0.385	0.567	

f-square - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON -> LY	0.047	0.052	0.026	1.844	0.065	
IMG -> LY	0.073	0.077	0.031	2.366	0.018	
SQ -> IMG	0.275	0.283	0.066	4.151	0.000	
SQ -> LY	0.011	0.013	0.011	1.041	0.298	
SQ -> TRS	0.265	0.273	0.067	3.980	0.000	
TRS -> LY	0.069	0.072	0.028	2.419	0.016	
ECON x TRS -> LY	0.015	0.017	0.012	1.261	0.207	

ตัวอย่าง Moderated Mediation Model with Quadratic Effect



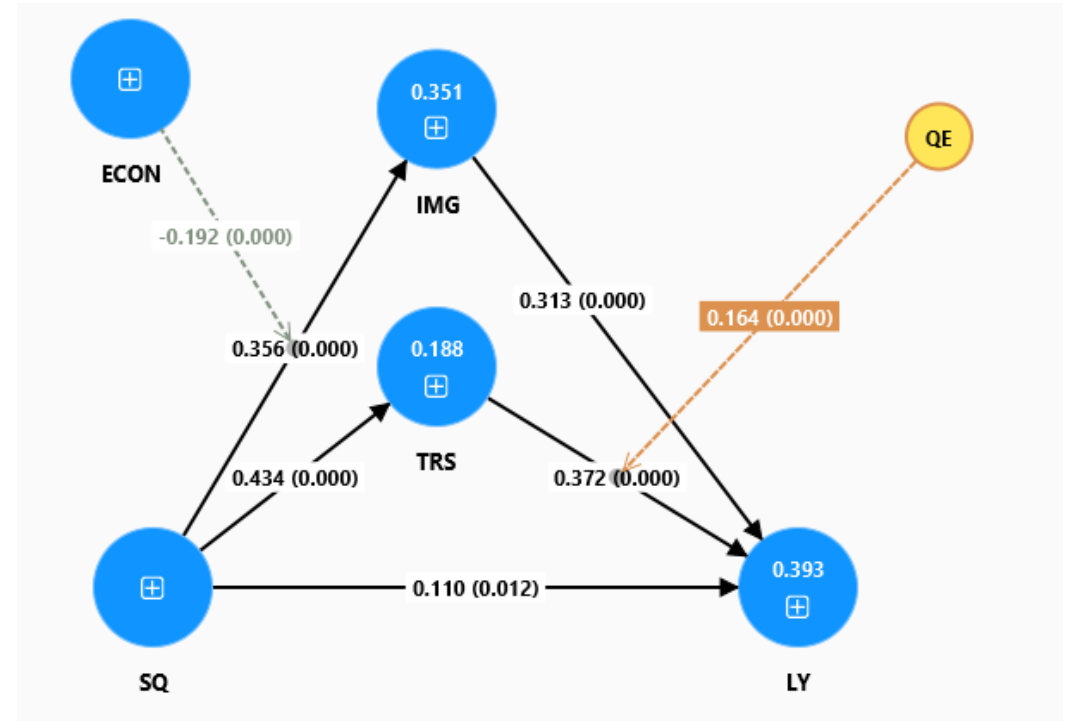
Moderated mediation model with quadratic effect

การวิเคราะห์การกำกับ พบว่าการกำกับของ Experience Economy (ECON) มีผลทางลบต่อความสัมพันธ์ระหว่าง SQ กับ IMG ตามเส้นทาง SQ->IMG แสดงว่าถ้า ECON มีค่าสูงขึ้นอิทธิพลที่ SQ มีต่อ IMG จะลดลงแสดงว่า ECON มีความสำคัญเหนือกว่า SQ ผู้ประกอบการด้านต่าง ๆ ในพื้นที่ปลายทางต้องให้ความสำคัญของ ECON มากขึ้นและควรทำ SQ ตามปกติ ไม่ต้องเน้น กลับกัน ถ้า ECON มีค่าลดลงหรือมีค่าน้อยหรือทำกันน้อย SQ จะมีอิทธิพลต่อ IMG มากขึ้น แสดงว่าในบรรยากาศที่ ECON ไม่ค่อยเด่น ผู้ประกอบการยังไม่ค่อยดำเนินกิจกรรมที่ทำให้ลูกค้ามีสิ่งจดจำ (memorable) และดื่มด่ำใจ (immersive) ลูกค้าไม่ค่อยทำรีวิวผ่าน SM กันมากนัก ในกรณีนี้ผู้ประกอบการในพื้นที่ปลายทางควรเน้น SQ ให้มาก เพื่อให้มีภาพลักษณ์ที่ดี และยังพบว่า TRS มีบทบาทการเชื่อมโยงสูงกว่า IMG

การวิเคราะห์ Quadratic Effect (QE) ผลการวิเคราะห์พบว่า QE มีอิทธิพลทางบวกและมีนัยสำคัญ แสดงว่า TRS (Tourist Satisfaction) มีได้มีอิทธิพลต่อความภักดีพื้นที่ปลายทาง (Destination Loyalty: LY) แบบเป็นเส้นตรง แต่กลับเป็นแบบเส้นโค้ง คือ Quadratic form) ดังสมการ

$$LY = 0.372 TRS^2 + 0.164 TRS$$

พบว่า $a = 0.372 > 0$ โค้งหงาย มีจุดเปลี่ยนโค้งที่ $TRS = \frac{-b}{2a} = \frac{-0.164}{2 \times 0.372} = -0.22$ ซึ่งเจตทอ้งโค้งคือ $(-0.22, -0.018)$



กราฟของสมการ $y = 372x^2 + 0.164x$

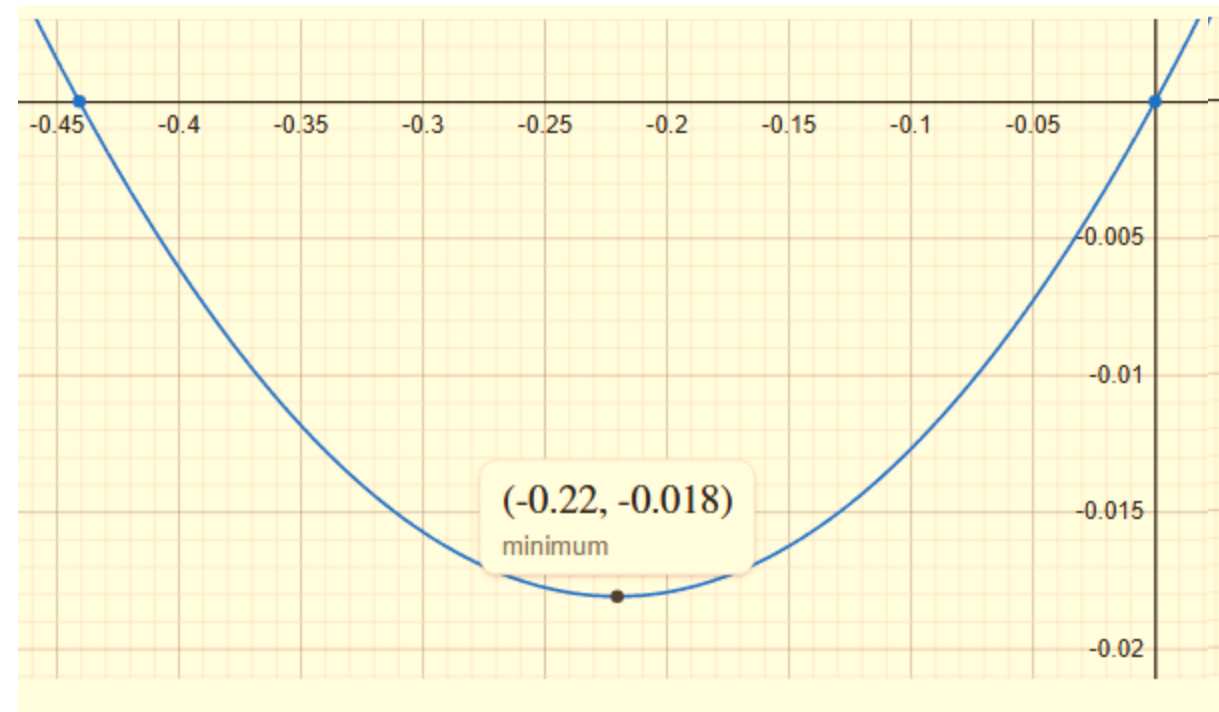
จากสมการ $y = ax^2 + bx + c$ จุดเปลี่ยนโค้งคือ $x = \frac{-b}{2a}$, ค่า y คำนวณได้จากการแทนที่ค่า x

ถ้า $a > 0$ ค่าต่ำสุดของ x จะอยู่ที่ท้องโค้งของ U-shape curve บ่งชี้ว่าเมื่อเพิ่มค่า x ไปค่าของ y จะสูงขึ้นอย่างรวดเร็วแบบอัตราเร่ง

ถ้า $a < 0$ ค่าสูงสุดของ x จะอยู่ที่ยอดโค้งของ Inverse U-shape curve บ่งชี้ว่านี่คือจุดสูงสุดแล้วต่อไปค่า y จะลดลงอย่างรวดเร็วแบบอัตราเร่ง

หมายเหตุ TRS มีค่าเฉลี่ยเท่ากับ 3.19, SD = 1.29 นั่นคือ

$$Z = \frac{X - 3.19}{1.29} = -0.22 \text{ ดังนั้น } X = 2.91$$

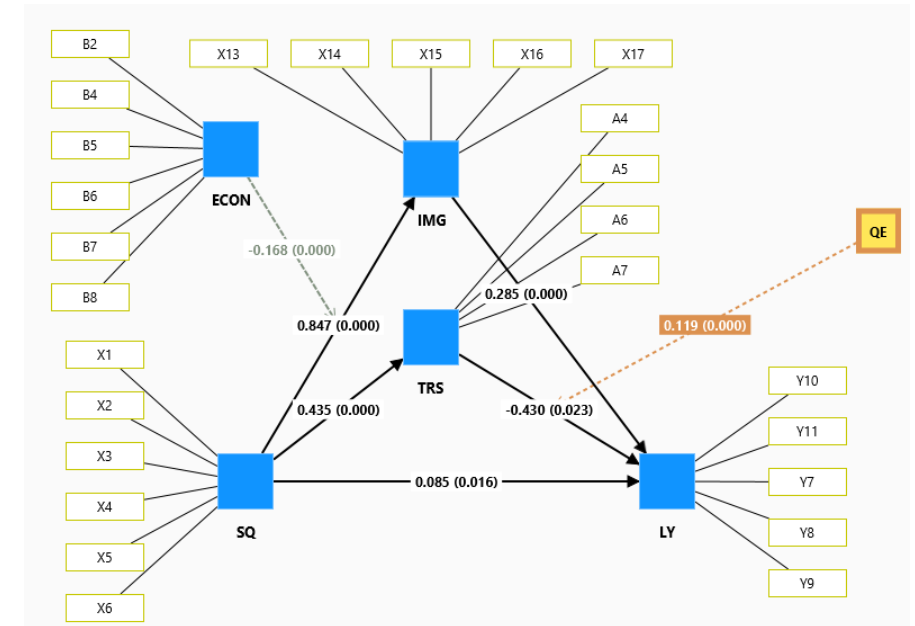


การวิเคราะห์ conditional Effect

เพื่อที่จะวิเคราะห์ Conditional Direct Effect (CDE) และ Conditional Indirect Effect (CIE) เราจะต้อง convert คือแปลงจาก PLS-SEM เป็น PROCESS แล้วสร้าง Bootstrapping ข้อมูลที่ใช้คือ construct score

หมายเหตุ

1. PROCESS ใช้วิเคราะห์ตัวแบบเส้นทาง (path model) ตัวแปรจะเป็น single-valued variable ดังนั้นถ้าตัวแปรเดิมเป็น multi-valued variable จะต้องแปลงข้อมูลเป็น construct score แล้ววิเคราะห์ตามวิธี PA โดย PROCESS จะมีแผ่นแบบ (template) ให้เลือกใช้ 92 ชุดและสามารถกำหนดเองเพิ่มได้ (customize) แต่ก็ต้องยึดเกณฑ์ของ PROCESS คือต้องมีตัวแปรสาเหตุและตัวแปรผลลัพธ์อย่างละ 1 ตัว ถูกสกรีนไปทางเดียว ภาพขวามือคือตัวแบบ PLS-SEM ที่แปลงมาเป็น path model แต่จะแสดงร่องรอยของตัวชี้วัดที่เป็นที่มาของ construct score
2. ใน PLS-SEM เราออกแบบกรอบแนวความคิดได้อิสระกว่า ไม่ยึดเกณฑ์ของ PROCESS ไม่ต้องรันตามแผ่นแบบ
3. ผลการวิเคราะห์ข้อมูลจาก PROCESS option ของ Smart PLS4 กับ PLS-SEM algorithm จะไม่ค่อยตรงกันเนื่องจากใช้อัลกอริทึมต่างกัน ดังนั้นการวิเคราะห์ตัวแบบโครงสร้างและตัวแบบมาตรวัดให้ใช้ตาม PLS-SEM algorithm ส่วนผลการวิเคราะห์ CDE และ CIE ให้ใช้ตาม PROCESS option ของ Smart PLS4



สรุปผลการวิเคราะห์ Moderated mediation model with quadratic effect

จากภาพและตารางสถิติสามารถสรุปได้ดังนี้

1. มาตรวัดมีความตรงและความเที่ยงสอดคล้องกับเกณฑ์ที่กำหนด
2. ทุกเส้นทางในกรอบแนวความคิดการวิจัยมีนัยสำคัญ อิทธิพลรวม อิทธิพลทางตรง อิทธิพลทางอ้อม อิทธิพลทางตรงอย่างมีเงื่อนไข อิทธิพลทางอ้อมอย่างมีเงื่อนไขและอิทธิพลกำลัง 2 ล้วนมีนัยสำคัญ

TE = 0.383, DE = 0.110,

IE1: (SQ → IMG → LY) = 0.111, IE2: (SQ → TRS → LY) = 0.161,

(TE, IE1 และ IE2 ใช้ผลการวิเคราะห์จาก PLS-SEM algorithm)

CIE1: (SQ → IMG → LY เมื่อ ECON = mean -1SD) = 0.128

CIE1: (SQ → IMG → LY เมื่อ ECON = mean) = 0.083

CIE1: (SQ → IMG → LY เมื่อ ECON = mean+1SD) = 0.038

(CIE1, CIE2 และ CIE3 ใช้ผลการวิเคราะห์จาก PROCESS option ของ Smart PLS4)

- 1) TE ลดลงจาก 0.383 เป็น 0.111 คือลดลง 71% แสดงว่า IMG และ TRS เป็นปัจจัยซ่อนเร้นที่มีอิทธิพลสูงมาก

สรุปผลการวิเคราะห์ Moderated mediation model with quadratic effect

2) อิทธิพลทางอ้อม IE1: (เส้นทาง $SQ \rightarrow IMG \rightarrow LY$) = 0.111, IE2: (เส้นทาง $SQ \rightarrow TRS \rightarrow LY$) = 0.161 มีนัยสำคัญทั้ง 2 โซ่เส้นทาง แต่ IE2 มีค่าสูงกว่า แสดงว่าทั้ง IMG และ TRS เป็นปัจจัยคนกลางที่มีความสำคัญแต่ TRS มีพลังสูงกว่าจึงต้องเน้น TRS มากกว่า

3) อิทธิพลทางอ้อมอย่างมีเงื่อนไขมีค่าสูงเมื่อ TRS มีค่าต่ำและค่อย ๆ สูงขึ้นเมื่อ ECON มีค่าสูงขึ้น หมายความว่าในบรรยากาศที่ ECON ยังไม่แผ่ขยายกว้างขวาง SQ จะมีอิทธิพลต่อ LY มากกว่าแสดงว่าต้องใส่ใจด้าน SQ ให้มากกว่า ECON ต่อมาเมื่อมีการปฏิบัติด้าน ECON มากขึ้น ๆ อิทธิพลของ SQ จะลดลงแสดงว่าต้องปฏิบัติด้าน ECON ให้มากขึ้น ส่วน SQ ก็ยังคงให้ความสำคัญ

4) QE มีนัยสำคัญและเป็นปริมาณบวกแสดงว่าในระยะแรก ๆ นักท่องเที่ยวอาจสนใจอย่างอื่นเช่น ความปลอดภัย ภาพลักษณ์ ความคุ้นเคย การมีราคาสินค้าและบริการที่ยุติธรรม เพิ่งมาเที่ยวครั้งแรกจึงยังอยู่ในระหว่างประเมิน/เปรียบเทียบ/เรียนรู้ เมื่อนานเข้าความพอใจจะค่อย ๆ สะสมจนถึงจุดหนึ่งอันเป็นจุดที่ความพึงพอใจมีค่าต่ำสุด (ในที่นี้คือ $X = 2.91$) จากนั้นไปถ้าพื้นที่ปลายทางสร้างความพึงพอใจให้นักท่องเที่ยวเพิ่มขึ้นเพียงเล็กน้อยความภักดีของนักท่องเที่ยวจะเพิ่มอย่างมากในลักษณะอัตราเร่ง

ผลการวิเคราะห์ข้อมูลจากการสำรวจจาก PROCESS option ของ Smart PLS4

Conditional direct effects - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
SQ -> IMG conditional on ECON at +1 SD	0.134	0.133	0.057	2.365	0.018	
SQ -> IMG conditional on ECON at -1 SD	0.447	0.447	0.050	9.025	0.000	
SQ -> IMG conditional on ECON at Mean	0.290	0.290	0.039	7.522	0.000	

Conditional indirect effects - Mean, STDEV, T values, p values							Copy t
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values		
SQ -> IMG -> LY conditional on ECON at +1 SD	0.038	0.038	0.018	2.173	0.030		
SQ -> IMG -> LY conditional on ECON at -1 SD	0.128	0.127	0.027	4.752	0.000		
SQ -> IMG -> LY conditional on ECON at Mean	0.083	0.083	0.018	4.498	0.000		

ผลการวิเคราะห์ข้อมูลจากการสำรวจจาก PLS-SEM algorithm

Specific indirect effects - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON -> IMG -> LY	0.098	0.100	0.022	4.520	0.000	
SQ -> TRS -> LY	0.162	0.162	0.029	5.666	0.000	
SQ -> IMG -> LY	0.111	0.112	0.023	4.789	0.000	
ECON x SQ -> IMG -> LY	-0.060	-0.060	0.017	3.483	0.001	

Total effects - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON -> IMG	0.314	0.317	0.045	6.959	0.000	
ECON -> LY	0.098	0.100	0.022	4.520	0.000	
IMG -> LY	0.313	0.315	0.049	6.346	0.000	
SQ -> IMG	0.356	0.356	0.044	8.173	0.000	
SQ -> LY	0.383	0.383	0.044	8.624	0.000	
SQ -> TRS	0.434	0.435	0.046	9.449	0.000	
TRS -> LY	0.372	0.373	0.050	7.439	0.000	
QE (TRS) -> LY	0.164	0.163	0.040	4.137	0.000	
ECON x SQ -> IMG	-0.192	-0.191	0.044	4.345	0.000	
ECON x SQ -> LY	-0.060	-0.060	0.017	3.483	0.001	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Heterotrait-monotrait ratio (HTMT) - Confidence intervals					
	Original sample (O)	Sample mean (M)	2.5%	97.5%	
IMG <-> ECON	0.517	0.517	0.415	0.612	
LY <-> ECON	0.539	0.538	0.443	0.627	
LY <-> IMG	0.611	0.611	0.527	0.693	
SQ <-> ECON	0.343	0.342	0.235	0.443	
SQ <-> IMG	0.533	0.531	0.432	0.623	
SQ <-> LY	0.419	0.418	0.317	0.512	
TRS <-> ECON	0.521	0.520	0.430	0.606	
TRS <-> IMG	0.604	0.604	0.518	0.684	
TRS <-> LY	0.602	0.602	0.525	0.676	
TRS <-> SQ	0.466	0.465	0.366	0.560	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

f-square - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON -> IMG	0.136	0.144	0.046	2.966	0.003	
IMG -> LY	0.107	0.113	0.038	2.850	0.004	
SQ -> IMG	0.175	0.180	0.049	3.560	0.000	
SQ -> LY	0.014	0.017	0.012	1.194	0.233	
SQ -> TRS	0.232	0.240	0.062	3.721	0.000	
TRS -> LY	0.145	0.148	0.039	3.673	0.000	
QE (TRS) -> LY	0.039	0.041	0.019	2.041	0.041	
ECON x SQ -> IMG	0.057	0.061	0.028	2.014	0.044	

Average variance extracted (AVE) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.703	0.703	0.018	39.421	0.000	
IMG	0.550	0.550	0.020	27.505	0.000	
LY	0.649	0.649	0.016	39.990	0.000	
SQ	0.782	0.782	0.013	61.973	0.000	
TRS	0.774	0.774	0.017	45.795	0.000	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Composite reliability (rho_a) - Mean, STDEV, T values, p values						
	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.920	0.922	0.007	128.511	0.000	
IMG	0.810	0.812	0.018	44.733	0.000	
LY	0.888	0.889	0.011	80.502	0.000	
SQ	0.949	0.950	0.004	222.423	0.000	
TRS	0.906	0.907	0.009	97.632	0.000	
QE (TRS)	1.000	1.000	0.000	n/a	n/a	
ECON x SQ	1.000	1.000	0.000	n/a	n/a	

ผลการวิเคราะห์ข้อมูลจากการสำรวจ

Cronbach's alpha - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.915	0.915	0.007	126.854	0.000	
IMG	0.790	0.790	0.018	42.876	0.000	
LY	0.855	0.855	0.012	73.970	0.000	
SQ	0.944	0.944	0.004	225.126	0.000	
TRS	0.903	0.902	0.009	95.657	0.000	

Composite reliability (rho_c) - Mean, STDEV, T values, p values

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values	
ECON	0.934	0.934	0.005	176.199	0.000	
IMG	0.857	0.856	0.011	78.942	0.000	
LY	0.899	0.899	0.007	128.201	0.000	
SQ	0.955	0.955	0.003	300.494	0.000	
TRS	0.932	0.932	0.006	151.287	0.000	

ภาคผนวก

Algorithm ของ PLS

1. Initial Assignment:

Set all outer weights to either 1 or $\frac{1}{\sqrt{n}}$ (n = number of indicators)

2. Outer Approximation:

Calculate outer estimates (Y) of latent variables (LVs) $Y = Xw$ where X = indicators matrix of size $k \times m$, w = weights of $m \times 1$, k = number of cases, m = indicators in a block, $Y = k \times 1$ estimated LV

3. Inner Approximation:

Calculate inner estimates (Z) using path relationships $Z = YE$ where E = inner weights 1×1 matrix

4. Weight Update: For reflective: $w = (X'X)(X'Z)(Z'Z)^{-1}$, For formative: $w = (\widehat{X'X})^{-1}X'Z$

5. Standardization:

Standardize weights and LV scores

6. Convergence Check:

Compare new vs old weights

Stop if difference < threshold (e.g., 10^{-5})

If not converged, return to step 2

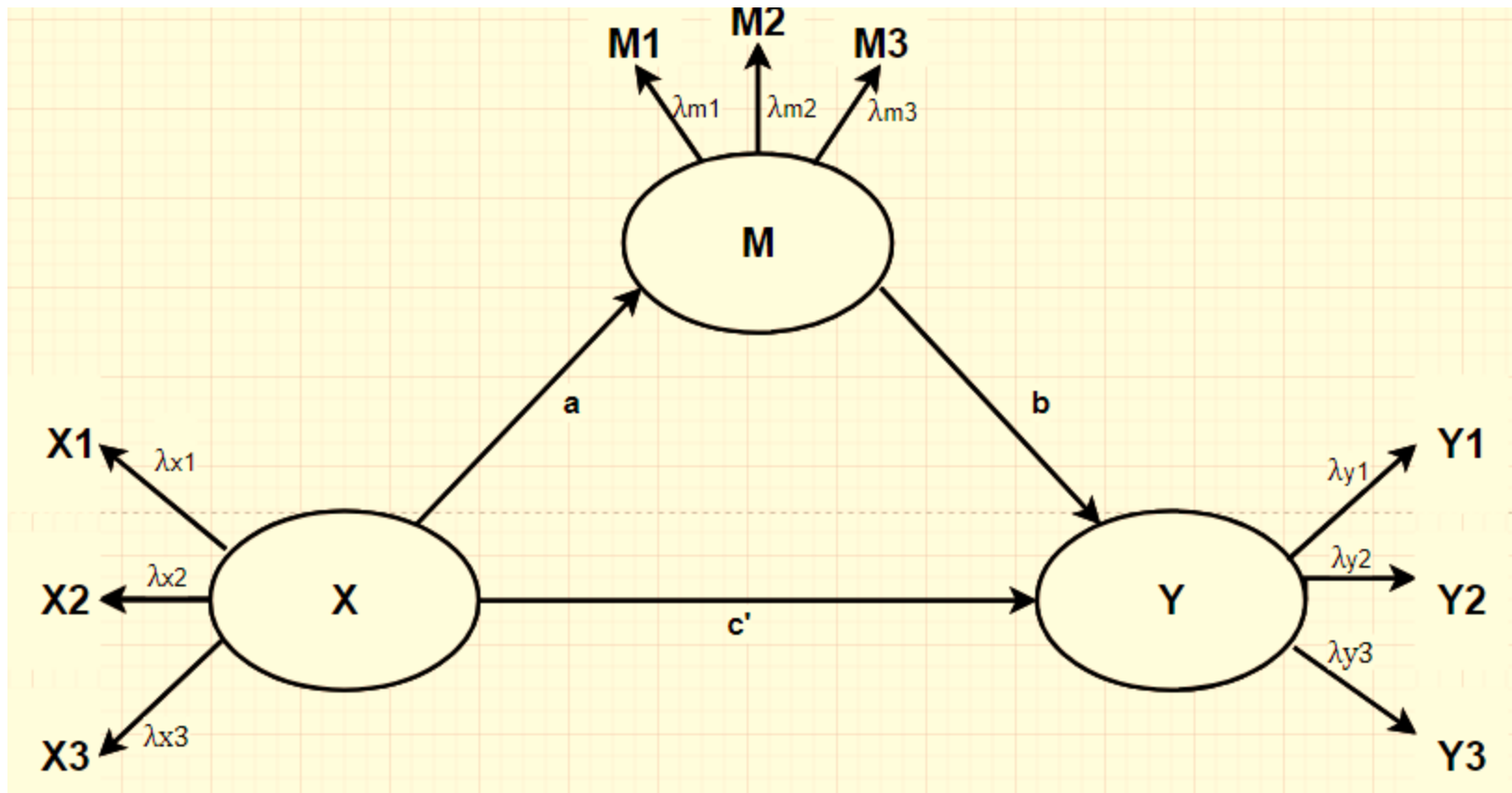
Note. PLS การวิเคราะห์ที่ละบล็อก

Algorithm ของ PLS อย่างง่าย

สมมุติตัวแปรอิสระคือ X ตัวแปรตามคือ Y ตัวแปรคันกลางคือ M แต่ละตัวแปรมีตัวชี้วัด 3 ตัว จะแสดง PLS algorithm โดยใช้ correlation สำหรับ outer model และใช้ regression analysis สำหรับ inner model โดยจะวนเวียนทำซ้ำจนค่า loading เสถียรโดยตัวชี้วัดมีข้อมูลสมมุติ 100 case ($n = 100$) แสดงให้เห็นโครงสร้างไฟล์เฉพาะ 5 case แรกคือ

PLS Algorithm แบบละเอียดและเข้าใจง่ายต่อไปนี้ถ้าเข้าใจดีจะสามารถเข้าใจอัลกอริทึมในสไลด์ก่อนหน้านี้ได้ง่ายขึ้น การประมาณค่าตัวแบบจะวนทำซ้ำเป็นรอบ ๆ (iteration) เริ่มจาก outer model ไปหา inner model และจาก inner model กลับไปหา outer model สลับไปมาจนกว่าค่าของน้ำหนักปัจจัยคงที่

Algorithm ของ PLS อย่างง่าย



Algorithm ของ PLS อย่างง่าย

Case	x1	x2	x3	m1	m2	m3	y1	y2	y3
1	3	4	2	5	3	4	2	3	4
2	4	2	3	3	4	2	4	2	3
3	2	3	4	4	2	3	3	4	2
4	5	3	4	2	3	4	5	3	4
5	4	5	3	3	5	4	4	5	3

Algorithm ของ PLS อย่างง่าย

1. standardize ข้อมูลของตัวชี้วัดทุกตัว ทำครั้งเดียวและใช้ในทุกรอบ
2. การ initialize น้ำหนักปัจจัย ครั้งแรกจะกำหนดให้น้ำหนักปัจจัย/สหสัมพันธ์เท่ากับ 1 ในทุกบล็อก
3. ประมวลค่าของตัวแปรแฝงด้วย construct score ซึ่งก็คือ weighted average ของคะแนนตัวชี้วัด ($\text{score} = \frac{\sum_j^k w_j x_j}{\sum_j^k w_j}$) โดย w_j คือค่าสหสัมพันธ์
4. Standardize ข้อมูลของ construct score ก่อนนำมาวิเคราะห์สมการถดถอยพหุที่กำหนดไว้ตามกรอบ จะได้ค่า standardized regression coefficient
5. predict ค่า construct score ที่เป็น dependent variable ในกรอบ ในที่นี้คือ M และ Y จะได้ค่าใหม่ของ construct score ให้ standardize ก่อนแล้วใช้ standardized construct score ของ M และ Y นี้ไปคำนวณหาค่าสหสัมพันธ์กับคะแนนตัวชี้วัดเดิมของตนจะได้น้ำหนักปัจจัยค่าใหม่
- ส่วนตัวแปรสาเหตุในที่นี้คือ X ให้ใช้ standardized construct score ของ X ไปคำนวณหาค่าสหสัมพันธ์กับคะแนนตัวชี้วัดเดิมจะได้ น้ำหนักปัจจัยค่าใหม่ได้เลย ไม่ต้องมีการ predict คะแนนปัจจัยของ X
6. วนซ้ำข้อ 3 ถึง 5 จนกระทั่งค่าน้ำหนักปัจจัยคงที่ (ไม่ใช่ค่าสัมประสิทธิ์ถดถอยคงที่เพราะน้ำหนักปัจจัยคงที่ค่าของสัมประสิทธิ์ถดถอยก็จะคงที่ด้วย) ปกติก็จะยอมให้คลาดเคลื่อนกับค่าที่ได้ในรอบก่อนหน้าประมาณ 0.001 คือ $\frac{1}{10^3}$ ถึง 0.0000001 คือ $\frac{1}{10^6}$

ตัวอย่าง PLS Algorithm อย่างง่าย

1. Standardize ข้อมูลของตัวชี้วัดได้ standardized indicator score ดังนี้

z_{x1} : -0.53, 0.35, -1.40, 1.23, 0.35

z_{x2} : 0.53, -1.23, -0.35, -0.35, 1.40

z_{x3} : -1.43, -0.24, 0.95, 0.95, -0.24

z_{m1} : 1.40, -0.35, 0.53, -1.23, -0.35

z_{m2} : -0.35, 0.53, -1.23, -0.35, 1.40

z_{m3} : 0.67, -1.57, -0.45, 0.67, 0.67

z_{y1} : -1.40, 0.35, -0.53, 1.23, 0.35

z_{y2} : -0.35, -1.23, 0.53, -0.35, 1.40

z_{y3} : 0.95, -0.24, -1.43, 0.95, -0.24

Algorithm ของ PLS อย่างง่าย

รอบที่ 1

1. กำหนดให้ค่าเริ่มต้นของน้ำหนักปัจจัย (initialized loading) มีค่า = 1 ทุกตัว

$$\mathbf{X}: w_{x1}=1, w_{x2}=1, w_{x3}=1$$

$$\mathbf{M}: w_{m1}=1, w_{m2}=1, w_{m3}=1$$

$$\mathbf{Y}: w_{y1}=1, w_{y2}=1, w_{y3}=1$$

2. คำนวณ Construct Scores รอบแรก:

$$\mathbf{X}_{\text{raw}} = (z_{x1} + z_{x2} + z_{x3})/3$$

$$\mathbf{M}_{\text{raw}} = (z_{m1} + z_{m2} + z_{m3})/3$$

$$\mathbf{Y}_{\text{raw}} = (z_{y1} + z_{y2} + z_{y3})/3$$

Algorithm ของ PLS อย่างง่าย

3. Standardize Construct Scores:

X_score: -0.48, -0.37, -0.27, 0.61, 0.50

M_score: 0.57, -0.46, -0.38, -0.30, 0.57

Y_score: -0.27, -0.37, -0.48, 0.61, 0.50

4. วิเคราะห์ MRA ระหว่าง standardized construct scores:

$$\mathbf{M_score} = 0.45(\mathbf{X_score}) + e$$

$$\mathbf{Y_score} = 0.30(\mathbf{X_score}) + 0.40(\mathbf{M_score}) + e$$

5. Predict ค่า construct score ของ M และ Y จากสมการถดถอยและ standardize ค่าของ construct score

6. หา correlations ระหว่าง standardized construct scores กับ standardized indicator scores

rx1,rx2,rx3: 0.75, 0.70, 0.68

rm1,rm2,rm3: 0.72, 0.75, 0.70

ry1,ry2,ry3: 0.70, 0.73, 0.71

Algorithm ของ PLS อย่างง่าย

รอบที่ 2

1. ใช้ correlations เป็น weights ใหม่หาค่า construct scores ของ X, M, Y ด้วยน้ำหนักใหม่

$$X_{\text{raw}} = (z_{x1} * 0.75 + z_{x2} * 0.70 + z_{x3} * 0.68) / (0.75 + 0.70 + 0.68)$$

$$M_{\text{raw}} = (z_{m1} * 0.72 + z_{m2} * 0.75 + z_{m3} * 0.70) / (0.72 + 0.75 + 0.70)$$

$$Y_{\text{raw}} = (z_{y1} * 0.70 + z_{y2} * 0.73 + z_{y3} * 0.71) / (0.70 + 0.73 + 0.71)$$

2. Standardize ค่า construct scores ของ X, M, Y ใหม่:

$$X_{\text{score}}: -0.52, -0.41, -0.33, 0.67, 0.59$$

$$M_{\text{score}}: 0.63, -0.51, -0.42, -0.33, 0.63$$

$$Y_{\text{score}}: -0.31, -0.41, -0.52, 0.67, 0.57$$

Algorithm ของ PLS อย่างง่าย

3. วิเคราะห์ MRA ระหว่าง standardized construct scores ได้สมการถดถอยใหม่คือ

$$M_score = 0.48(X_score) + e$$

$$Y_score = 0.32(X_score) + 0.42(M_score) + e$$

แล้ว predict ค่า construct score ของ M และ Y แล้ว standardize ค่าของ X, M และ Y

หาค่า correlations ใหม่ระหว่าง standardized score กับตัวชี้วัด:

$$rx1, rx2, rx3: \quad 0.77, 0.73, 0.70$$

$$rm1, rm2, rm3: 0.74, 0.77, 0.72$$

$$ry1, ry2, ry3: \quad 0.72, 0.75, 0.73$$

Algorithm ของ PLS อย่างง่าย

รอบที่ 3

1. หาค่า construct scores ใหม่ใช้ นำหนักคือค่าสหสัมพันธ์ในรอบที่ 2 แล้วหา standardized construct scores:

X_score: -0.54, -0.43, -0.34, 0.70, 0.61

M_score: 0.65, -0.53, -0.44, -0.34, 0.66

Y_score: -0.32, -0.43, -0.54, 0.70, 0.59

ค่าเปลี่ยนแปลงประมาณ 0.1-0.2 จากรอบที่ 2

2. วิเคราะห์ MRA ระหว่าง standardized construct scores ได้สมการถดถอยใหม่คือ

$$\mathbf{M_score} = 0.50(\mathbf{X_score}) + e$$

$$\mathbf{Y_score} = 0.43(\mathbf{X_score}) + 0.33(\mathbf{M_score}) + e$$

predict ค่า construct score ของ M และ Y แล้ว standardize ค่าของ X, M และ Y

Algorithm ของ PLS อย่างง่าย

3. หาค่า correlations ใหม่ระหว่าง standardized score กับตัวชี้วัด:

$rx1, rx2, rx3:$ 0.78, 0.74, 0.71

$rm1, rm2, rm3:$ 0.75, 0.78, 0.73

$ry1, ry2, ry3:$ 0.73, 0.76, 0.74

รอบที่ 4

ค่าเปลี่ยนแปลงประมาณน้อยกว่า 0.001

ผลลัพธ์คือ

Algorithm ของ PLS อย่างง่าย

รอบที่ 5

$$\mathbf{M_score} = 0.50(\mathbf{X_score}) + \mathbf{e}; R^2 = 0.25$$

$$\mathbf{Y_score} = 0.43(\mathbf{X_score}) + 0.33(\mathbf{M_score}) + \mathbf{e}; R^2 = 0.36$$

Loadings:

$$\mathbf{X}: \lambda_{x1} = 0.78, \lambda_{x2} = 0.74, \lambda_{x3} = 0.71$$

$$\mathbf{M}: \lambda_{m1} = 0.75, \lambda_{m2} = 0.78, \lambda_{m3} = 0.73$$

$$\mathbf{Y}: \lambda_{y1} = 0.73, \lambda_{y2} = 0.76, \lambda_{y3} = 0.74$$

Algorithm ของ PLS อย่างง่าย

หมายเหตุ

1. Construct Score คือ weighted average
$$= \frac{\sum_j^k w_j x_j}{\sum_j^k w_j}$$
2. ค่าประมาณของสัมประสิทธิ์การถดถอยทุกค่าต้องเป็น standardized coefficient ดังนั้นจึงต้อง standardize ข้อมูลก่อนการวิเคราะห์ทุกครั้ง ควรทราบว่าค่าสหสัมพันธ์ระหว่าง standardized construct score กับ standardized indicator score ก็คือสัมประสิทธิ์ถดถอยซึ่งมีค่าเท่ากันเสมอถ้าคำนวณจาก standardized score
3. Convergent rule คือผลต่างของ loading (ก็คือ correlation หรือ simple regression coefficient) จากรอบก่อนหน้ากับรอบใหม่ต้องน้อยกว่า 1 ใน 1,000 (คือ 10^{-3}) ถึง 1 ใน 1,000,000 (คือ 10^{-6}) หรือน้อยกว่านี้ เราสามารถกำหนดค่าได้
4. เราจะแปลงค่าข้อมูลของตัวชี้วัดทุกตัวด้วยวิธี standardize เพียงครั้งเดียวในตอนเริ่มต้นและใช้ค่านั้นตลอดไปในทุกรอบ (iteration) แต่ construct scores เราจะต้อง standardize ใหม่ทุกรอบเพราะ construct score มีค่าที่คำนวณมาใหม่เสมอจากค่าใหม่ของสัมประสิทธิ์ถดถอยและค่าใหม่ของสหสัมพันธ์

CB-SEM Algorithm

1. **CFA: Compare S (observed) vs Σ (implied) matrices**
2. **Model Specification: $\Sigma(\theta) = \Lambda\Phi\Lambda' + \Theta$**
3. **Maximum Likelihood: $\text{Min } F = \log |\Sigma(\theta)| + \text{tr}(S\widehat{\Sigma(\theta)}^{-1}) - \log |S| - p$**
4. **Parameter Estimation: Iterative optimization using Newton-Raphson**
5. **Fit Assessment: $\chi^2 = (N-1)[F]$, $df = [p(p+1)/2] - q$**

Where: S = Sample covariance matrix, Σ = Model-implied covariance matrix, Λ = Factor loading matrix

Φ = Factor covariance matrix, Θ = Error covariance matrix, p = Number of observed variables

N = Sample size, q = Number of parameters, θ = Model parameters

The process iterates until convergence criterion is met (typically when change in $\chi^2 < \text{threshold}$)

References

มนตรี พิริยะกุล (2544), เทคนิคการวิเคราะห์สมการถดถอย, มหาวิทยาลัยรามคำแหง.

Bagozzi, R. P., & Yi, Y. (1988). On the evaluation of structural equation models. *Journal of the Academy of Marketing Science*, 16(1), 74-94.

Briggs, S. R., & Cheek, J. M. (1986). The role of factor analysis in the development and evaluation of personality scales. *Journal of Personality*, 54(1), 106-148.

Chin, W. W. (1998). The partial least squares approach structural equation modeling. *Modern Methods for Business Research*, 295(2), 295-336.

Clark, L. A., & Watson, D. (1995). Constructing validity: Basic issues in objective scale development. *Psychological Assessment*, 7(3), 309-319.

Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Lawrence Erlbaum Associates.

Dijkstra, T. K., & Henseler, J. (2015). Consistent partial least squares path modeling. *MIS Quarterly*, 39(2), 297-316.

References

- Draper, N. R., & Smith, H. (1998). Applied regression analysis (3rd ed.). Wiley.**
- Falk, R. F., & Miller, N. B. (1992). A primer for soft modeling. University of Akron Press.**
- Field, A. (2013). Discovering statistics using IBM SPSS statistics (4th ed.). Sage.**
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39-50.**
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2014). A primer on partial least squares structural equation modeling (PLS-SEM). Sage Publications.**
- Hair, J. F., Ringle, C. M., & Sarstedt, M. (2011). PLS-SEM: Indeed, a silver bullet. *Journal of Marketing Theory and Practice*, 19(2), 139-152.**
- Hair, J. F., Sarstedt, M., Ringle, C. M., & Gudergan, S. P. (2018). Advanced issues in partial least squares structural equation modeling (PLS-SEM). Sage.**

References

- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 1-21.**
- Henseler, J., Ringle, C. M., & Sinkovics, R. R. (2009). The use of partial least squares path modeling in international marketing. *Advances in International Marketing*, 20, 277-319.**
- Kline, R. B. (2011). Principles and practice of structural equation modeling (3rd ed.). Guilford Press.**
- Lohmöller, J. B. (1989). Latent variable path modeling with partial least squares. Physica-Verlag.**
- Lovie, P. (2005). Coefficient of variation. *Encyclopedia of Statistics in Behavioral Science*.**
- Moore, D. S., Notz, W. I, & Flinger, M. A. (2013). The basic practice of statistics (6th ed.). W. H. Freeman and Company.**

References

- Nunnally, J. C., & Bernstein, I. H. (1994). *Psychometric theory* (3rd ed.). McGraw-Hill.
- Rovinelli, R. J., & Hambleton, R. K. (1977). On the use of content specialists in the assessment of criterion-referenced test item validity. *Dutch Journal of Educational Research*, 2, 49-60.
- Turner, R. C., & Carlson, L. (2003). Indexes of item-objective congruence for multidimensional items. *International Journal of Testing*, 3(2), 163-171.
- Wetzels, M., Odekerken-Schröder, G., & Van Oppen, C. (2009). Using PLS path modeling for assessing hierarchical construct models: Guidelines and empirical illustration. *MIS Quarterly*, 33(1), 177-195.
- Bollen, K. A., & Lennox, R. (1991). Conventional wisdom on measurement: A structural equation perspective. *Psychological Bulletin*, 110(2), 305-314. <https://doi.org/10.1037/0033-2909.110.2.305>
- Diamantopoulos, A., & Siguaw, J. A. (2006). Formative versus reflective indicators in organizational measure development: A comparison and empirical illustration. *British Journal of Management*, 17(4), 263-282. <https://doi.org/10.1111/j.1467-8551.2006.00500.x>
- Diamantopoulos, A., & Winklhofer, H. M. (2001). Index construction with formative indicators: An alternative to scale development. *Journal of Marketing Research*, 38(2), 269-277. <https://doi.org/10.1509/jmkr.38.2.269.18845>
- Diamantopoulos, A., Riefler, P., & Roth, K. P. (2008). Advancing formative measurement models. *Journal of Business Research*, 61(12), 1203-1218. <https://doi.org/10.1016/j.jbusres.2008.01.009>

References

- Edwards, J. R., & Bagozzi, R. P. (2000). On the nature and direction of relationships between constructs and measures. *Psychological Methods*, 5(2), 155-174. <https://doi.org/10.1037/1082-989X.5.2.155>
- Jarvis, C. B., MacKenzie, S. B., & Podsakoff, P. M. (2003). A critical review of construct indicators and measurement model misspecification in marketing and consumer research. *Journal of Consumer Research*, 30(2), 199-218. <https://doi.org/10.1086/376806>
- MacKenzie, S. B., Podsakoff, P. M., & Jarvis, C. B. (2005). The problem of measurement model misspecification in behavioral and organizational research and some recommended solutions. *Journal of Applied Psychology*, 90(4), 710-730. <https://doi.org/10.1037/0021-9010.90.4.710>
- MacKenzie, S. B., Podsakoff, P. M., & Podsakoff, N. P. (2011). Construct measurement and validation procedures in MIS and behavioral research: Integrating new and existing techniques. *MIS Quarterly*, 35(2), 293-334. <https://doi.org/10.2307/23044045>
- Petter, S., Straub, D., & Rai, A. (2007). Specifying formative constructs in information systems research. *MIS Quarterly*, 31(4), 623-656. <https://doi.org/10.2307/25148814>