

ข้อได้เปรียบของโมเดลเชิงเส้นระดับลดหลั่นของ Raudenbush & Bryk (2002)

ฉบับปรับปรุงครั้งที่ 2

Advantages of Hierarchical Linear Models (HLM) of Raudenbush & Bryk

(2002) Second Edition

เพ็ญภัคร พินผา¹

บทคัดย่อ

โมเดลเชิงเส้นระดับลดหลั่นของ Raudenbush & Bryk (2002) ฉบับปรับปรุงครั้งที่ 2 มีข้อได้เปรียบมากกว่าฉบับปรับปรุงครั้งที่ 1 ของ Bryk & Raudenbush (1992) ดังนี้ 1) โครงสร้างข้อมูลระดับลดหลั่น HLM ฉบับปรับปรุงครั้งที่ 2 มีความคล้ายคลึงกับข้อมูลที่มีอยู่ในสังคมหรือในธรรมชาติ 2) มีการปรับปรุงการประมาณค่าของอิทธิพลระดับบุคคล 3) มีโมเดลอิทธิพลข้ามระดับ 4) สามารถแยกส่วนประกอบความแปรปรวน-ความแปรปรวนร่วมได้ 5) มีการขยายขอบเขตของตัวแปรตามมากขึ้น 6) สามารถแยกส่วนประกอบความแปรปรวน – ความแปรปรวนร่วมได้ 7) สามารถวิเคราะห์ข้อมูลที่มีโครงสร้างแบบข้ามกลุ่มหรือข้ามหน่วยในระดับเดียวกันได้ 8) สามารถวิเคราะห์โมเดลตัวแปรแฝงได้ และสุดท้าย 9) การอ้างอิงแบบเบย์เซียน

คำสำคัญ : โมเดลเชิงเส้นระดับลดหลั่น, โมเดลเชิงเส้นพหุระดับ, โมเดลอิทธิพลข้ามระดับ

Abstract

The Advantages of Hierarchical Linear Models (HLM) of Raudenbush & Bryk (2002) second edition when compare with the first edition of Bryk & Raudenbush (1992) are 1) hierarchical data structure second edition is familiar with a common phenomenon 2) improved estimation of individual effects 3) modeling cross-level effects 4) partitioning variance-covariance components 5) an expanded range of outcome variables 6) incorporating cross-classified data structures 7) multivariate models 8) latent variables models and 9) Bayesian inference.

Keywords : Hierarchical Linear Model, Multi-level Linear Model, Cross-level Model

¹ อาจารย์ ดร. คณะบริหารศาสตร์ มหาวิทยาลัยอุบลราชธานี

ผู้เขียนได้มีโอกาสทำวิจัยและเรียนวิชา Hierarchical Linear Models (HLM) ของ Raudenbush & Bryk (2002) ฉบับปรับปรุงครั้งที่ 2 กับ Professor Dr. Mike Seltzer ที่ภาควิชา Social Research Methodology ที่มหาวิทยาลัยแคลิฟอร์เนีย แห้งนครลอสแอนเจลิส (UCLA) ซึ่งเป็นผู้วิจัยวิธีการเบย์เซียน (Bayesian Approaches) ในบทที่ 13 ในตำราเล่มนี้ ของ Raudenbush & Bryk (2002) Seltzer กล่าวว่าตำราเล่มนี้มีความสมบูรณ์แบบกว่าฉบับปรับปรุงครั้งที่ 1 และเป็นตำราที่เป็นที่ยอมรับกันทั่วโลก นอกจากนี้ยังไม่นับนักวิชาการหรือนักสถิติท่านใดเขียนตำราหรือคิดค้นโปรแกรมการวิเคราะห์ข้อมูลทุกระดับเชิงเส้นได้ดีกว่าโปรแกรม HLM ของ Raudenbush & Bryk (2002) ในขณะนี้ ตำราเล่มนี้ใช้มาถึง 12 ปีแล้ว (ค.ศ. 2002 – 2014) ซึ่ง Raudenbush & Bryk ก็ยังไม่มีแผนในการปรับปรุงหนังสือเล่มนี้แต่อย่างใด ผู้เขียนจึงได้ทำการวิเคราะห์โมเดลเชิงเส้นระดับลดหลั่นที่ Raudenbush & Bryk ได้พัฒนาขึ้นนี้มาอย่างดีและเป็นยอมรับในแวดวงวิธีวิทยาการวิจัยทางสังคมศาสตร์ การศึกษา จิตวิทยา สาธารณสุขหรือทางการแพทย์นั้นน่าจะประกอบด้วยเหตุผลดังนี้

1) โครงสร้างข้อมูลระดับลดหลั่น HLM มีความคล้ายคลึงกับข้อมูลที่มีอยู่ในสังคมหรือในธรรมชาติ

งานวิจัยทางสังคมศาสตร์โดยส่วนใหญ่มีความเกี่ยวข้องกับข้อมูลที่มีโครงสร้างระดับลดหลั่น ในองค์กรที่นักวิจัยทำการศึกษาก็พบว่าลักษณะขององค์กรหรือสังคมของมนุษย์นั้นจะประกอบไปด้วยกลุ่ม หรือหน่วยงานต่างๆ เช่น ครอบครัว โรงเรียน มหาวิทยาลัย โรงพยาบาล บริษัท ชุมชน ตำบล อำเภอ จังหวัด และประเทศ เป็นต้น ในองค์กรต่างๆ เหล่านี้ก็ล้วนแล้วแต่มีหัวหน้า รองหัวหน้า หัวหน้าแผนก ผู้ปฏิบัติงานหรือสมาชิกในองค์กร เป็นต้น การดำเนินงานโดยส่วนใหญ่ก็มักมีคำสั่งมาจากหัวหน้ากลุ่มหรือหัวหน้าองค์กรเป็นนโยบายหรือแนวปฏิบัติก่อนเรื่อยไปถึงระดับล่างสุดคือระดับบุคคล การตัดสินใจก็จะมาจากหน่วยงานส่วนบนก่อนเช่น นโยบายจากระดับเขตพื้นที่การศึกษา ย่อมมีอิทธิพลต่อโรงเรียน ครูผู้ปฏิบัติงาน และนักเรียน ตามลำดับ ทั้งเขตพื้นที่การศึกษา โรงเรียน ครู และนักเรียน ถือเป็นหน่วยในการวิเคราะห์ที่แตกต่างกัน และสามารถวิเคราะห์ได้ถึงสี่ระดับ ตัวอย่างงานวิจัยของ Kyriakides (2004) ได้ทำวิจัยเรื่องความแตกต่างของประสิทธิผลของโรงเรียนที่มีส่วนสัมพันธ์กับเพศและสังคมในห้องเรียน โดยแบ่งการวิเคราะห์ข้อมูลออกเป็น 2 ระดับคือ ระดับโรงเรียน จำนวน 58 โรงเรียน และระดับนักเรียน เป็นนักเรียนชั้นประถมศึกษาชั้นปีที่ 2 จำนวน 1,778 คน ใน Cyprus ประเทศเนเธอร์แลนด์ เก็บข้อมูลโดยใช้แบบสอบถามโดยมีค่าสัมประสิทธิ์ความเที่ยง Cronbach Alpha เท่ากับ 0.82 ผลการวิจัยพบว่า เพศ สังคมในห้องเรียน และการเข้าชั้นเรียนมีความสัมพันธ์กับพัฒนาการทางด้านคณิตศาสตร์ของนักเรียน แต่เมื่อวิเคราะห์แบบทุกระดับกับไม่พบว่า เพศ และสังคมในห้องเรียนทำให้ประสิทธิผลของโรงเรียนแตกต่างกันอย่างมีนัยสำคัญทางสถิติ Levacic และ Jenkins (2006) ได้ประเมินประสิทธิผลของโรงเรียนลักษณะพิเศษในประเทศอังกฤษ โดยใช้การวิเคราะห์ทุกระดับ กับกลุ่มตัวอย่างที่เป็นโรงเรียนพิเศษกับโรงเรียนปกติในระดับมัธยมศึกษาปี 2001

ตัวแปรควบคุมได้แก่ ความรู้เดิมของเด็กโดยใช้คะแนนในช่วงชั้นที่ 2 เพศ อายุ และบริบทของโรงเรียน ประสิทธิภาพของโรงเรียนแต่ละโรงเรียนได้จากคะแนนมูลค่าเพิ่ม (value-added) ที่ได้จากการวัดใน โมเดลพหุระดับ 2 ระดับ คือ ระดับนักเรียนและระดับโรงเรียน ผลการวิจัยพบว่า ประสิทธิภาพของ โรงเรียนแตกต่างกันตามชนิดของโรงเรียน (โรงเรียนรัฐบาลหรือเอกชน) และระยะเวลาที่โรงเรียน เหล่านั้นเป็นโรงเรียนพิเศษหรือดีเด่น Mason และคณะ (1983) ได้ทำการศึกษาปัจจัยด้านเศรษฐกิจ และสังคมระหว่างประเทศที่มีส่วนสัมพันธ์กับผลสัมฤทธิ์ทางการเรียนของนักเรียนและอัตราการเกิดของ คนในประเทศต่างๆ แตกต่างกันหรือไม่ ผลการวิจัยพบว่า ความแตกต่างทางเศรษฐกิจของแต่ละ ประเทศ โดยดูจากผลิตภัณฑ์มวลรวมภายในประเทศ (Gross National Product, GNP) จาก 15 ประเทศ ที่อาศัยอยู่ในเมืองหรือชนบทมีปฏิสัมพันธ์กับผลสัมฤทธิ์ทางการศึกษาของนักเรียนและมี อิทธิพลต่ออัตราการเกิดของคนในประเทศด้วยกลุ่มครอบครัวที่พ่อแม่มีการศึกษาสูงอัตราการเกิดใน ครอบครัวนั้นจะต่ำกว่าครอบครัวที่พ่อแม่มีการศึกษาต่ำ

จากงานวิจัยดังกล่าวสังเกตได้ว่า สิ่งที่น่าสนใจศึกษาโดยเฉพาะบุคคลต้องเกี่ยวข้องกับ กลุ่มสังคมหรือองค์กรเสมอจากเล็กไปใหญ่ (เพราะมนุษย์เป็นสัตว์สังคม) เช่น นักเรียนเป็นบุคคลที่อยู่ใน องค์กรครอบครัว โรงเรียน กลุ่มเพื่อน กลุ่มเพื่อนร่วมชั้นเรียน และอาศัยอยู่ในตำบล อำเภอ จังหวัด และประเทศ ลักษณะของข้อมูลแบบนี้แสดงให้เห็นถึงโครงสร้างของข้อมูลระดับลดหลั่น ซึ่งมีมนุษย์ทุก คนล้วนแล้วแต่อยู่ภายใต้โครงสร้างขององค์กร (Osborne, 2000) อย่างเป็นระดับลดหลั่นอย่าง หลีกเลียงไม่ได้

นอกจากงานวิจัยข้างต้นแล้ว ข้อมูลแบบพัฒนาการหรือการติดตามการเปลี่ยนแปลงตาม ช่วงเวลา การเก็บรวบรวมข้อมูลซ้ำๆ ในช่วงเวลาต่างกัน ข้อมูลเหล่านี้เป็นข้อมูลที่มีโครงสร้างสอดแทรก (Nested structure) อยู่ภายในตัวบุคคล และบุคคลก็อยู่ในกลุ่มชั้นเรียน ในโรงเรียน ตัวอย่างงานวิจัย เช่น von Hippel (2009) ได้ศึกษาประสิทธิภาพของโรงเรียนด้วยพิจารณาจากข้อมูล 3 ลักษณะคือ 1) ระดับผลสัมฤทธิ์ทางการเรียน 2) จากอัตราการเรียนรู้อื่น และ 3) จากอิทธิพลของช่วงเวลา (seasonal “impact”) โดยใช้ข้อมูลระยะยาวจากเด็กอนุบาลปี 1998-99 จาก The Only National US Survey แบ่งการประมาณค่าอัตราการเรียนรู้อื่นเป็นช่วงเวลาภาคฤดูร้อนกับภาคปกติ กลุ่มตัวอย่างเป็นโรงเรียน จำนวน 287 โรง วิเคราะห์ข้อมูลโดยใช้โมเดลพัฒนาการพหุระดับ (multilevel growth model) ของ Raudenbush และ Bryk (2002) โมเดลการวิเคราะห์มี 3 ระดับ คือ ระดับการทดสอบ (level 1) ภายในนักเรียนแต่ละคน ระดับนักเรียน (level 2) และระดับโรงเรียน (level 3) ผลการวิเคราะห์ ปรากฏว่า ความตรงจากการวัดประสิทธิภาพของโรงเรียนด้วยอัตราการเรียนรู้อื่นและอิทธิพลของช่วงเวลามี ค่ามากกว่าความตรงจากการวัดประสิทธิภาพของโรงเรียนโดยใช้ผลสัมฤทธิ์ทางการเรียน นอกจากนี้ยังมี งานวิจัยมากมายนับไม่ถ้วนที่ใช้โมเดลการวิเคราะห์แบบโมเดลเชิงเส้นระดับลดหลั่น (HLM) ทั้งในและ ต่างประเทศ อาทิ Choi (2011) ; Secker และ Lissitz (1997); Abbott และคณะ (2002); Kyriakides

และคณะ (2009); จตุภูมิ เขตจัตุรัส (2552); เพ็ญภัคร พันธ์ผา (2554) และศุภลักษณ์ ใจแสงทรัพย์ (2547) เป็นต้น

สรุป โมเดลเชิงเส้นระดับลดหลั่น (HLM) มีความเหมาะสมในการวิเคราะห์ข้อมูลทางสังคมศาสตร์ที่มีลักษณะเป็นโครงสร้างระดับลดหลั่น (hierarchical) หรือสอดแทรก (nested) กันได้อย่างเหมาะสมเพราะ 1) HLM สามารถวิเคราะห์หน่วย (Units) ที่เราสนใจศึกษาทั้งที่มีเงื่อนไขเหมือนกันหรือแตกต่างกันได้ 2) HLM สามารถวิเคราะห์ข้อมูลที่เราสนใจศึกษาได้มากกว่า 1 ระดับ ที่สอดแทรกอยู่ในอีกระดับได้ และ 3) HLM สามารถวิเคราะห์ข้อมูลช่วงเวลาที่แตกต่างกันแม้จะเก็บภายใต้หน่วยเดียวกัน

2) โปรแกรม HLM ฉบับปรับปรุงครั้งที่ 2 มีการปรับปรุงการประมาณค่าของอิทธิพลระดับบุคคล (Improved Estimation of Individual Effects)

ในการวิเคราะห์ข้อมูลโดยทั่วไป นักวิจัยมักจะกังวลเกี่ยวกับการใช้คะแนนทดสอบมาตรฐาน (standardized test scores) สำหรับกลุ่มตัวอย่างที่มีขนาดเล็ก เช่น การศึกษานักเรียนชนกลุ่มน้อยในแต่ละโรงเรียน เนื่องจากนักเรียนชนกลุ่มน้อยในแต่ละโรงเรียนมีน้อยจึงทำให้ข้อมูลไม่เพียงพอในการวิเคราะห์ในสมการ และการแยกวิเคราะห์เฉพาะนักเรียนชนกลุ่มน้อยออกจากขนาดตัวอย่างเล็กแล้ว วิธีการยังยุ่งยากเพราะต้องแยกนักเรียนชนกลุ่มน้อยออกจากกลุ่มนักเรียนในทุกโรงเรียน การที่ข้อมูลมีขนาดตัวอย่างที่เพียงพอถือว่ามีส่วนสำคัญมากในการเป็นตัวแทนทำนายค่าพารามิเตอร์ของประชากรเพราะจะทำให้สัมประสิทธิ์การวิเคราะห์ไม่ลำเอียง การใช้โปรแกรม HLM จึงถือว่าสามารถช่วยแก้ปัญหานี้ได้ โดย HLM จะทำการนำข้อมูลทั้งหมดมาเป็นข้อมูลประมาณค่าตัวแปรทำนายของแต่ละโรงเรียน เช่น ตัวแปรเชื้อชาติของนักเรียนที่แบ่งเป็น ขนผิวขาวกับชนกลุ่มน้อย ตัวแปรทำนายนี้ในแต่ละโรงเรียนจะถูกถ่วงน้ำหนักด้วยจำนวนข้อมูลจากโรงเรียนนั้นๆ การถ่วงน้ำหนักจากจำนวนของตัวแปรนั้นก็จะทำให้ค่าที่ประมาณได้มีความแม่นยำมากขึ้น (Morris, 1983)

3) HLM มีโมเดลอิทธิพลข้ามระดับ (Modeling Cross-Level Effects)

โมเดลระดับลดหลั่นนั้นสร้างขึ้นเพื่อทดสอบสมมติฐานที่ว่าตัวแปรที่วัดในระดับที่ 1 มีความสัมพันธ์กับตัวแปรระดับที่สูงกว่าอย่างไร เพราะอิทธิพลข้ามระดับมีอยู่ในงานวิจัยทางพฤติกรรมศาสตร์และสังคมศาสตร์ กรอบแนวคิดนี้จึงมีความสำคัญกว่าแนวคิดในการวิเคราะห์แบบดั้งเดิม ตัวอย่างงานวิจัยทางสังคมศาสตร์ Huttenlocher และคณะ (1991) ได้เก็บรวบรวมข้อมูลระยะยาว (Longitudinal data) จากเด็กทั้งหมด 7 ครั้ง อายุตั้งแต่ 14 – 26 เดือน เพื่อดูพัฒนาการทางด้านคำศัพท์ของเด็กแต่ละคน ผลปรากฏว่าพัฒนาการทางคำศัพท์ของเด็กขึ้นอยู่กับเพศและการพูดของมารดาและสภาพแวดล้อมในบ้านที่เป็นข้อมูลต่างระดับกัน ซึ่งการวิเคราะห์ด้วยโปรแกรม HLM นี้ผลการวิเคราะห์พัฒนาการทางด้านคำศัพท์ของเด็กที่ตัวแปรมาจากสภาพแวดล้อมในบ้านพบว่ามีความก้าวหน้าวิจัยที่ผ่านมาที่ใช้การวิเคราะห์แบบดั้งเดิม ซึ่งจากการวิเคราะห์โมเดลเชิงเส้นระดับลดหลั่นนี้

สามารถแสดงให้เห็นถึงความสัมพันธ์ข้ามระดับที่มีความซับซ้อนได้โดยพัฒนาแนวคิดมาจากการวิเคราะห์แบบดั้งเดิม (conventional approaches) คือการวิเคราะห์ multiple regression ดังนั้นโมเดล HLM จึงเป็นโมเดลที่สามารถวิเคราะห์อิทธิพลข้ามระดับได้จึงเป็นข้อได้เปรียบข้อโมเดล HLM นี้ Unsupervised

4) โปรแกรม HLM สามารถแยกส่วนประกอบความแปรปรวน – ความแปรปรวนร่วม (Partitioning Variance-Covariance Component) ได้

โมเดลเชิงเส้นระดับลดหลั่น (HLM) สามารถประมาณค่าสัดส่วนความแปรปรวน-ความแปรปรวนร่วมสำหรับกลุ่มข้อมูลที่มีจำนวนไม่เท่ากันและมีหลายระดับ โดยแบ่งการประมาณค่าออกเป็น อิทธิพลคงที่ (fixed effects) และอิทธิพลสุ่ม (random effects) ในรูป Variance-Covariance เมทริกซ์ สัดส่วนความแปรปรวนที่ HLM ประมาณค่าจะสามารถบอกถึงอิทธิพลของตัวแปรทำนายที่มีอยู่ในโมเดล HLM และสามารถบอกตัวแปรตามมีความผันแปรเท่าไรในแต่ละระดับ ตัวอย่างงานวิจัยของ เพ็ญภักดิ์ พันผา (2554) ที่ได้พัฒนาโมเดลมูลค่าเพิ่มพหุระดับ 2 ระดับ เพื่อการวัดประสิทธิผลของโรงเรียน โดยเก็บข้อมูลจากนักเรียนในระดับชั้นมัธยมศึกษาปีที่ 3 จำนวน 1,852 คน จาก 49 โรงเรียน ที่เกี่ยวกับผลสัมฤทธิ์ทางการเรียนและตัวแปรที่ไม่เกี่ยวข้องกับโรงเรียน (non-school variables) และตัวแปรที่เป็นการดำเนินงานของโรงเรียน ผลการวิจัยปรากฏว่า ตัวแปรที่ไม่เกี่ยวข้องกับโรงเรียนด้านคุณลักษณะของผู้เรียน เช่น ผลสัมฤทธิ์เดิม โอกาสในการเรียนพิเศษ เพศ ความคาดหวัง ส่งผลต่อสัมฤทธิ์ทางการเรียนของนักเรียน นอกจากนี้ในระดับโรงเรียน ผลสัมฤทธิ์ทางการเรียนของนักเรียนยังขึ้นอยู่กับคุณภาพการสอนของครูกับขนาดของโรงเรียนอีกด้วย โดยตัวแปรทำนายทั้ง 2 ระดับ สามารถอธิบายความแปรปรวนของตัวแปรตาม (ผลสัมฤทธิ์ทางการเรียนของนักเรียน) ได้ถึง 70% ส่วนเมทริกซ์ความแปรปรวนร่วมหรือสหสัมพันธ์ภายในชั้น (intra-class correlation) ที่ได้มีค่าเท่ากับ 0.4817 แสดงว่า 48.17% เป็นความแปรปรวนของผลสัมฤทธิ์ทางการเรียนระหว่างโรงเรียน และอีก 51.83% เป็นความแปรปรวนของผลสัมฤทธิ์ทางการเรียนระหว่างนักเรียน อย่างไรก็ตามการใช้โมเดล HLM วิเคราะห์เมทริกซ์ความแปรปรวน – ความแปรปรวนร่วมที่ได้ก็ไม่เที่ยงตรงเสมอไป ในกรณีที่โมเดลมีความซับซ้อนและมีหลายระดับ การแปลความหมายก็จะสับสนและมีข้อผิดพลาดเกิดขึ้น เพราะนักวิจัยจะต้องนำค่าความคลาดเคลื่อนที่ได้ในแต่ละระดับมาคำนวณหาสหสัมพันธ์ภายในชั้นเอง

5) HLM มีการขยายขอบเขตของตัวแปรตาม (An Expanded Range of Outcome Variables)

ในหนังสือ HLM ฉบับปรับปรุงครั้งที่ 1 ตัวแปรตามในระดับที่ 1 (outcomes at level 1) ต้องเป็น ตัวแปรต่อเนื่องเท่านั้น เพราะมีข้อตกลงเบื้องต้นว่าความคลาดเคลื่อนที่เกิดขึ้นจะต้องมีการแจกแจงปกติ แต่โมเดล HLM ฉบับปรับปรุงครั้งที่ 2 กับตัวแปรตามสามารถใช้ได้ทั้งตัวแปรแบบจัดประเภท (discrete outcomes) โดยเฉพาะตัวแปรที่มีการแจกแจงแบบทวินาม (เช่น การถูกจ้างงาน-

ไม่ถูกจ้างงาน พึ่งพอใจในงาน-ไม่พึงพอใจในงาน ต่ำ-ปานกลาง-สูง เป็นต้น) การแจกแจงแบบ Logistic ซึ่งตัวแปรตามเหล่านี้การแจกแจงของข้อมูลไม่ใช่การแจกแจงแบบปกติและไม่เป็นเชิงเส้น การที่ตัวแปรตามมีการแจกแจงไม่ปกติ (non-normal outcomes) ทำให้มีนักวิจัยหลายๆ ท่าน ในสมัยนั้นพัฒนาโปรแกรมหรือสถิติในการคำนวณให้เหมาะสมกับข้อมูล อาทิ Stiratelli, Laird และ Ware (1984) Wong และ Mason (1985) (อ้างถึงใน Raudenbush และ Bryk, 2002) ได้แก้ปัญหานี้โดยใช้การประมาณค่าลำดับที่ 1 (first order) แบบ maximum likelihood (ML) ต่อมา Goldstein และ Longford (1993) (อ้างถึงใน Raudenbush และ Bryk, 2002) ได้พัฒนาซอฟต์แวร์เพื่อช่วยประมาณค่าโมเดลที่ตัวแปรตามเป็นแบบจัดประเภท 2-3 โมเดล แต่อย่างไรก็ตาม Breslow และ Clayton (1993) German และ Rodriguez (1995) (อ้างถึงใน Raudenbush และ Bryk, 2002) ได้แสดงให้เห็นว่าการประมาณค่าดังกล่าวไม่ค่อยแม่นยำในบางเงื่อนไข ดังนั้น Goldstein (1995) ได้พัฒนาการประมาณลำดับที่ 2 (second order approximation) ขึ้นมา Hedeker และ Gibbons (1993) และ Pinheiro และ Bates (1995) (อ้างถึงใน Raudenbush และ Bryk, 2002) ได้พัฒนาวิธีประมาณค่าแบบ Gauss-Hermite quadrature ที่มีความแม่นยำกว่า ML และใช้ในโปรแกรม SAS Proc Mixed และโปรแกรม Mixor ส่วนวิธีการที่มีความแม่นยำและสะดวกในการคำนวณ Raudenbush, Yang และ Yosef (2000) ได้ใช้การเปลี่ยนแปลงแบบลาเพลสขั้นสูง (high-order Laplace transform) ในโปรแกรม HLM ผลของวิวัฒนาการทางด้านการสถิติการคำนวณทำให้นักวิจัยสามารถประยุกต์ใช้โมเดล Logic และตัวแปรตามที่อยู่ในมาตรวัดต่าง ๆ ในโปรแกรม HLM ได้

6) โปรแกรม HLM สามารถวิเคราะห์ข้อมูลที่มีโครงสร้างแบบข้ามกลุ่มหรือข้ามหน่วยในระดับเดียวกัน (Incorporating Cross-Classified Data Structures) ได้

HLM ในฉบับปรับปรุงที่ 1 ถ้าหากนักวิจัยจะวิเคราะห์ข้อมูลพหุระดับ ข้อมูลจะต้องเป็นระดับลดหลั่นอย่างสมบูรณ์ (pure hierarchies) คือ นักเรียนต้องอยู่ภายในห้องเรียนใดห้องเรียนหนึ่ง และห้องเรียนก็ต้องอยู่ภายในโรงเรียนใดโรงเรียนหนึ่งเท่านั้น และเมื่อมีการปรับปรุงโปรแกรม HLM ทำให้สามารถวิเคราะห์ข้อมูลสอดแทรกได้ซับซ้อนมากยิ่งขึ้น ยกตัวอย่างเช่น การศึกษาผลสัมฤทธิ์ทางการเรียนของนักเรียนที่ได้เข้าเรียนในโรงเรียนใดโรงเรียนหนึ่งและนักเรียนก็อยู่ในกลุ่มเพื่อนบ้านกลุ่มใดกลุ่มหนึ่ง ข้อมูลที่ได้ไม่มีจำเป็นต้องเป็นข้อมูลระดับลดหลั่นอย่างสมบูรณ์ เพราะโดยหลักแล้วนักเรียนก็มาจากกลุ่มเพื่อนบ้านและขณะเดียวกันกลุ่มเพื่อนบ้านก็ส่งบุตรหลานเข้าโรงเรียน ซึ่งการศึกษาแบบนี้นักเรียนจะสามารถอยู่ได้ทั้งสองกลุ่มคืออยู่ในกลุ่มเพื่อนบ้านและโรงเรียนซึ่งเป็นหน่วยในระดับเดียวกัน ซึ่งสูตรและโมเดลการทดสอบข้อมูลลักษณะนี้แบบไม่มีเงื่อนไขเป็นการทดสอบอิทธิพลร่วมในกรณีแบบข้ามกลุ่ม ยกตัวอย่างเช่น การวิเคราะห์หาความผันแปร (variance) ของนักเรียนที่อยู่ภายในโรงเรียนและภายในกลุ่มเพื่อนบ้าน และระหว่างโรงเรียนและระหว่างกลุ่มเพื่อนบ้านที่อยู่ในระดับเดียวกัน โมเดลแบบไม่มีเงื่อนไขที่ไม่มีตัวทำนายใด ๆ ในโมเดลการทดสอบแบบข้ามกลุ่มจะดังนี้

โมเดลแบบไม่มีเงื่อนไข (Unconditional Model)

ระดับที่ 1 โมเดลภายในเซลล์หรือภายในกลุ่ม (within – cell model)

ในโมเดลนี้ นักเรียนจะสอดแทรกอยู่ทั้งกลุ่มเพื่อนบ้านและโรงเรียน ดังนั้นโมเดลจะเป็นดังนี้

$$Y_{ijk} = \pi_{0jk} + e_{ijk}, \quad e_{ijk} \sim N(0, \sigma^2)$$

เมื่อ Y_{ijk} หมายถึง ผลสัมฤทธิ์ของนักเรียนคนที่ i อยู่ในกลุ่มเพื่อนบ้านที่ j และอยู่ในโรงเรียนที่ k

π_{0jk} หมายถึง ค่าเฉลี่ยผลสัมฤทธิ์ทางการเรียนของนักเรียนที่อยู่ในกลุ่มหรือเซลล์ (cell) jk

ซึ่งหมายความว่านักเรียนคนนั้นอยู่ในกลุ่มเพื่อนบ้านที่ j และอยู่ในโรงเรียนที่ k

e_{ijk} หมายถึง ค่าอิทธิพลสุ่ม (อิทธิพลนักเรียน) ที่คะแนนของนักเรียนคนที่ ijk คลาดเคลื่อน

หรือเบี่ยงเบน (deviation) ไปจากค่าเฉลี่ยของกลุ่มหรือเซลล์ (cell mean)

โดยค่าคลาดเคลื่อนที่ได้ ยอมรับ (assume) ว่ามีการแจกแจงปกติมีค่าเฉลี่ยเป็น

ศูนย์และความแปรปรวนมีค่าเท่ากับ σ^2

ระดับที่ 2 โมเดลระหว่างเซลล์หรือระหว่างกลุ่ม (between – cell model)

ผลการวิเคราะห์ในโมเดลนี้ ความผันแปรที่ได้คือ อิทธิพลจากกลุ่มเพื่อนบ้าน อิทธิพลของโรงเรียน และอิทธิพลปฏิสัมพันธ์ระหว่างกลุ่มเพื่อนบ้านและโรงเรียน โมเดลในระดับที่ 2 เป็นดังนี้

$$\pi_{0jk} = \theta_0 + b_{00j} + c_{00k} + d_{0jk}$$

$$b_{00j} \sim N(0, \tau_{b00}),$$

$$c_{00k} \sim N(0, \tau_{c00}),$$

$$d_{0jk} \sim N(0, \tau_{d00}),$$

เมื่อ θ_0 คือ ค่าเฉลี่ยรวมผลสัมฤทธิ์ทางการเรียนของนักเรียนทุกคน (grand mean)

b_{00j} คือ ค่าอิทธิพลสุ่มหลักที่มาจากกลุ่มเพื่อนบ้านที่ j ซึ่งหมายความว่าค่าเฉลี่ยที่เกิดจากกลุ่มเพื่อนบ้านที่ j โดยเฉลี่ยจากทุกโรงเรียน (averaged over all schools) ค่าที่ได้มีการแจกแจงปกติมีค่าเฉลี่ยเป็นศูนย์และความแปรปรวนมีค่าเท่ากับ τ_{b00}

c_{00k} คือ ค่าอิทธิพลสุ่มหลักที่มาจากโรงเรียนที่ k ซึ่งหมายความว่าค่าเฉลี่ยที่เกิดจากโรงเรียนที่ k โดยเฉลี่ยจากทุกกลุ่มเพื่อนบ้าน (averaged over all neighborhoods) ค่าที่ได้มีการแจกแจงปกติมีค่าเฉลี่ยเป็นศูนย์และความแปรปรวนมีค่าเท่ากับ τ_{c00}

d_{00k} คือ ค่าอิทธิพลปฏิสัมพันธ์ระหว่างเพื่อนบ้านที่ j โรงเรียนที่ k ความคลาดเคลื่อนหรือส่วนเบี่ยงเบนที่เกิดขึ้นมาจากค่าเฉลี่ยของกลุ่ม (cell mean) หักออกจากค่าเฉลี่ยรวม (grand mean) และค่าอิทธิพลหลัก 2 ค่า (two main effects) ค่าที่ได้มีการแจกแจงปกติมีค่าเฉลี่ยเป็นศูนย์และความแปรปรวนมีค่าเท่ากับ τ_{b00} ดังนั้น จะได้โมเดลรวม (combined model) ดังนี้

$$Y_{ijk} = \theta_0 + b_{00j} + c_{00k} + d_{0jk} + e_{ijk},$$

ส่วนการหาค่าสัดส่วนความแปรปรวนและค่าความเที่ยงของโมเดลการทดสอบแบบไม่มีเงื่อนไขแบบข้ามกลุ่มเป็นดังนี้ (อ้างถึงใน Raudenbush และ Bryk, 2002)

- 1) การหาสหสัมพันธ์ระหว่างผลสัมฤทธิ์ทางการเรียนของนักเรียนสองคนที่อาศัยอยู่ในกลุ่มเพื่อนบ้านเดียวกันและโรงเรียนเดียวกันสูตรการหาสหสัมพันธ์เป็นดังนี้

$$\text{Corr}(Y_{ijk}, Y_{i'jk}) = \rho_{bcd} = \frac{\tau_{b00} + \tau_{c00} + \tau_{d00}}{\tau_{b00} + \tau_{c00} + \tau_{d00} + \sigma^2}$$

- 2) การหาสหสัมพันธ์ระหว่างผลสัมฤทธิ์ทางการเรียนของนักเรียนสองคนที่อาศัยอยู่ในกลุ่มเพื่อนบ้านเดียวกันแต่อยู่คนละโรงเรียนสูตรการหาสหสัมพันธ์เป็นดังนี้

$$\text{Corr}(Y_{ijk}, Y_{i'jk'}) = \rho_b = \frac{\tau_{b00}}{\tau_{b00} + \tau_{c00} + \tau_{d00} + \sigma^2}$$

- 3) การหาสหสัมพันธ์ระหว่างผลสัมฤทธิ์ทางการเรียนของนักเรียนสองคนที่อาศัยอยู่ต่างกลุ่มเพื่อนบ้านแต่อยู่ในโรงเรียนเดียวกันสูตรการหาสหสัมพันธ์เป็นดังนี้

$$\text{Corr}(Y_{ijk}, Y_{i'jk}) = \rho_c = \frac{\tau_{c00}}{\tau_{b00} + \tau_{c00} + \tau_{d00} + \sigma^2}$$

การหาค่าความเที่ยงเมื่ออิทธิพลของเพื่อนบ้านสามารถจำแนกนักเรียนในการเข้าเรียนในโรงเรียนแต่ละโรงได้ดังนี้

$$\text{Reliability}[(\hat{b}_{00j} + \hat{d}_{0jk}) | c_{00k}] = \frac{\tau_{b00} + \tau_{c00} + \tau_{d00}}{\tau_{b00} + \tau_{c00} + \tau_{d00} + \sigma^2}$$

การหาความเที่ยงเมื่ออิทธิพลของโรงเรียนสามารถจำแนกนักเรียนที่มาจากกลุ่มเพื่อนบ้านเดียวกันมีสูตรคำนวณดังนี้

$$\text{Reliability}[(\hat{c}_{00k} + \hat{d}_{0jk}) | b_{00j}] = \frac{\tau_{c00} + \tau_{d00}}{\tau_{c00} + \tau_{d00} + \sigma^2 / n_{jk}}$$

สรุป โมเดล HLM ฉบับปรับปรุงที่ 2 สามารถขยายข้อจำกัดในการวิเคราะห์ข้อมูลระดับลดหลั่นสมบูรณ์ (pure hierarchies) มาเป็นโมเดลการวิเคราะห์หือทธิพลสุ่มแบบข้ามกลุ่มได้ (models for cross-classified random effects) หากผู้อ่านสนใจในรูปแบบการวิเคราะห์นี้สามารถอ่านเพิ่มเติมได้ในบทที่ 12 ของ Raudenbush และ Bryk (2002)

7) โปรแกรม HLM สามารถวิเคราะห์โมเดลพหุตัวแปร (Multivariate Models) ได้

ใน HLM ฉบับปรับปรุงที่ 1 ตัวแปรตามในโมเดลการวิเคราะห์มีได้เพียง 1 ตัว และอยู่ในระดับที่ 1 ต่อมาได้มีการพัฒนาให้โปรแกรม HLM สามารถวิเคราะห์ตัวแปรตามได้มากกว่า 1 ตัว ที่เรียกว่าโมเดลพหุตัวแปร (Multivariate Models) โดยส่วนใหญ่การประยุกต์ใช้การวิเคราะห์พหุตัวแปรนี้จะอยู่ในรูปข้อมูลที่มีการวัดซ้ำ (repeated-measures data) หรือข้อมูลระยะยาว (longitudinal data) หรือข้อมูลอนุกรมเวลา (time-series data) โดยกำหนดให้ข้อมูลนั้นมาจากบุคคลเดียวกันในแต่ละวัดละช่วงเวลาที่แตกต่างกัน หรือการวิเคราะห์พัฒนาการ (growth models) ทั้งที่เป็นแบบเชิงเส้น (linear growth model) และแบบสมการกำลังสอง (quadratic growth model) หรือโมเดลพัฒนาการแบบอื่น ๆ

8) HLM สามารถวิเคราะห์โมเดลตัวแปรแฝง (Latent Variable Models) ได้ด้วย

ตัวแปรแฝงโดยส่วนใหญ่วัดได้จากแปรสังเกตต่างๆ เป็นตัวบ่งชี้และคำนวณในรูปแบบสหสัมพันธ์ระหว่างตัวแปรสังเกตได้ การวิเคราะห์ข้อมูลที่เป็นตัวแปรแฝงจะใช้วิธีการคล้ายกับการวัดความคลาดเคลื่อนที่ได้จากตัวแปรสังเกตได้ในโมเดลถดถอยพหุในระดับที่ 1 โมเดลเชิงเส้นระดับลดหลั่นจะหาความสัมพันธ์ระหว่างความคลาดเคลื่อนที่ได้ในตัวแปรสังเกตได้ (fallible observed data) กับตัวแปรแฝง เช่นความคลาดเคลื่อนในการวัดของตัวแปรคุณลักษณะของกลุ่มเช่น โรงเรียน กลุ่มเพื่อนบ้าน ต่อมา HLM จะทำการประมาณค่าอิทธิพลทางตรงและอิทธิพลทางอ้อมของความคลาดเคลื่อนที่ได้จากการวัดตัวแปรต้นในบริบทโมเดลเชิงเส้นระดับลดหลั่น นอกจากนี้โมเดลการตอบสนองข้อสอบ (item response models) ที่มีการใช้ในการทดสอบทางด้านการศึกษาก็เช่นเดียวกัน บางโมเดลสามารถนำเสนอฟังก์ชันความน่าจะเป็นคุณลักษณะการตอบสนองของข้อสอบ ที่บอกถึงความสามารถหรือคุณสมบัติแฝง (latent trait) ของบุคคล ซึ่งคุณลักษณะแฝงเป็นตัวแปรบุคคลที่เราสนใจศึกษา บางโมเดลก็สามารถนำเสนอโมเดลสองระดับเกี่ยวกับการตอบสนองของข้อสอบที่สอดแทรกอยู่ภายในบุคคล จากตัวอย่างที่กล่าวข้างต้น HLM จึงเป็นโปรแกรมที่สามารถวิเคราะห์ข้อมูลแฝงหรือตัวแปรแฝงที่มีลักษณะเป็นระดับลดหลั่น

9) HLM ใช้วิธีการอ้างอิงแบบเบย์เซียน (Bayesian Inference)

HLM ในเวอร์ชันที่ 1 การวิเคราะห์สถิติเชิงอ้างอิงทั้งหมดจะใช้วิธีความน่าจะเป็นสูงสุด (maximum likelihood, ML) ในการประมาณค่าพารามิเตอร์ ซึ่งวิธี ML นั้นค่าพารามิเตอร์ที่ต้องประมาณค่าจะต้องมีความสอดคล้อง (consistent) และเส้นโค้งในกราฟจะต้องไม่ลำเอียง

(asymptotically unbiased) และมีประสิทธิภาพ (efficient) การประมาณค่าโดยวิธีนี้จะต้องใช้กลุ่มตัวอย่างที่มีขนาดใหญ่และมีการแจกแจงปกติ จึงจะทำให้การทดสอบด้วยสถิติที่มีความไวในการทดสอบอย่างใดก็ตามเงื่อนไขของกลุ่มตัวอย่างต้องมีขนาดใหญ่ (large sample) ในกรณีที่เป็นการวิเคราะห์พหุระดับ หน่วยการวิเคราะห์ในระดับสูงสุดจะเป็นตัวกำหนดขนาดของกลุ่มตัวอย่างในระดับล่าง ดังนั้น ในกรณีที่จำนวนกลุ่มตัวอย่างในระดับสูงสุดมีน้อยสถิติเชิงอ้างอิงที่ใช้ในการทดสอบแบบ ML จะไม่ค่อยน่าเชื่อถือ ดังนั้น ความน่าเชื่อถือจากการวิเคราะห์ข้อมูลจึงขึ้นอยู่กับจำนวนข้อมูลหรือขนาดของกลุ่มตัวอย่างที่ใช้ วิธีการแบบเบย์เซียนจึงเป็นวิธีการใหม่ที่ยืดหยุ่นและสะดวกต่อการคำนวณ โดยเฉพาะในบริบทข้อมูลแบบพหุระดับ วิธีการนี้จะใช้ชุดคำสั่งตามขั้นตอน (algorithms) โดยวิธีการ Monte Carlo Method ในการประมาณค่าการแจกแจงภายหลัง (posterior distribution) ซึ่งในอดีตนั้นสามารถจัดการหรือทำได้ยาก (intractable) วิธีการนี้ยังรวมข้อมูลการเพิ่มขึ้น (data augmentation) (Tanner & Wong, 1987 อ้างถึงใน Raudenbush & Bryk, 2002) และ Gibbs Sampling (Gelfand และคณะ, 1990; Gelfand & Smith, 1990 อ้างถึงใน Raudenbush & Bryk, 2002) ซึ่งซอฟต์แวร์ HLM ใช้อยู่ในขณะนี้ จากการทดลองของ Seltzer (1993) ที่ได้วิเคราะห์ข้อมูลซ้ำจากข้อมูลของ Huttenlocher และคณะ (1991) เกี่ยวกับพัฒนาการทางด้านคำศัพท์ของนักเรียน 22 คน ในระหว่างที่ศึกษาในชั้นประถมศึกษาชั้นปีที่ 2 พบว่าวิธีการ ML มีความเสี่ยง (risk) เพราะจำนวนนักเรียนในระดับที่สองมีขนาดเล็ก แต่เมื่อใช้วิธีเบย์เซียนในการประมาณค่าแม้ว่ากลุ่มตัวอย่างในระดับสูงจะมีขนาดเล็กแต่ค่าสถิติที่ได้ค่อนข้างแข็งแกร่งถึงแม้กลุ่มตัวอย่างจะไม่เพียงพอก็ตาม ส่วนรายละเอียดวิธีการประมาณค่าแบบเบย์เซียนจะได้กล่าวถึงในบทความวิชาการถัดไปว่ามีรูปแบบการประมาณค่าอย่างไร และแตกต่างจากวิธีการดั้งเดิม Maximum Likelihood มากน้อยแค่ไหน

เอกสารอ้างอิง

- จตุภูมิ เขตจัตุรัส. การพัฒนาโมเดลมูลค่าเพิ่มของผลสัมฤทธิ์ทางวิชาการและแบบตรวจสอบรายการประเมินตนเองเพื่อเพิ่มมูลค่ากระบวนการจัดการศึกษา. วิทยานิพนธ์ปริญญาดุษฎีบัณฑิต, ภาควิชาวิจัยและจิตวิทยาการศึกษา บัณฑิตวิทยาลัย จุฬาลงกรณ์มหาวิทยาลัย, 2552.
- เพ็ญภัคร พันผา. การพัฒนาโมเดลมูลค่าเพิ่มในระดับเพื่อการวัดประสิทธิผลของโรงเรียน. วิทยานิพนธ์ปริญญาดุษฎีบัณฑิต, ภาควิชาวิจัยและจิตวิทยาการศึกษา บัณฑิตวิทยาลัย จุฬาลงกรณ์มหาวิทยาลัย, 2554.
- ศุภลักษณ์ ใจแสวงทรัพย์. ปัจจัยที่มีผลต่อคะแนนพัฒนาการวิชาคณิตศาสตร์ของนักเรียนชั้นมัธยมศึกษาปีที่ 3. วิทยานิพนธ์ปริญญามหาบัณฑิต, ภาควิชาวิจัยและจิตวิทยาการศึกษา บัณฑิตวิทยาลัย จุฬาลงกรณ์มหาวิทยาลัย, 2547.

- Abbott, L. M., Joireman, J. & Stroh, R. H. "The influence of district size, school size, and socioeconomic status on student achievement in Washington: a replication study using hierarchical linear modeling." Washington School Research Center. Technical Report#3. [Online]. Available from:<http://www.spu.edu/orgs/research/WSRC%20HLM%20District%20Size%20Final%2010-2-02.pdf> [2012, July, 4], 2002.
- Bryk, A.S., & Raudenbush, S. W. "Hierarchical linear models: Applications and data analysis methods." Newbury Park, CA: Sage Publications, 1992.
- Choi K. "Latent variable regression 4-level hierarchical model using multisite multiple-cohorts longitudinal data." CRESST REPORT 801. National Center for Research on Evaluation, Standard, & Student Testing. UCLA, Graduate School of Education & Information Studies, 2011.
- Goldschmidt, P., Martinez, J. F., Niemi, D., & Baker, E. L. "Relationships among measures as empirical evidence of validity: incorporating multiple indicators of achievement and school context," Educational assessment. 12(3&4), 239–266, 2007.
- Huttenlocher, J. E., Haight, W., Bryk, A. S., & Seltzer, M. "Early vocabulary growth: relation to language input and gender," Developmental Psychology. 27(2), 236–249, 1991.
- Kyriakides, L. "Differential school effectiveness in relation to sex and social class: some implications for policy evaluation," Educational Research and Evaluation. 10(2), 141 – 161, 2004.
- Kyriakides, L., Creemers, B. M. P., & Antonio, P. "Teacher behavior and student outcomes: Suggestions for research on teacher training and professional development," Teaching and Teacher Education. 1-12, 2009.
- Levacic, R., & Jenkins, A. "Evaluating the effectiveness of specialist school in England," School Effectiveness and School Improvement. 17(3), 229 – 254, 2006.
- Morris, C. N. "Parametric empirical bayes inference: theory and applications," Journal of the American Statistical Association. 78, 321- 344, 1983.

- Osborne, W. J. Advantages of hierarchical linear modeling. Practical Assessment, Research & Evaluation, 7(1). Retrieved August 2, 2012 from <http://PAREonline.net/getvn.asp?v=7&n=1>, 2000.
- Ruadenbush, S. W., Yang, M., & Yosef, M. “Maximum likelihood for hierarchical models via high-order, multivariate LaPlace approximation,” Journal of Computational and Graphical Statistics. 9(1), 141-157, 2000.
- Ruadenbush, S. W., & Bryk, A. S. “Hierarchical linear models: Applications and Data analysis methods. (second edition),” California: sage Publications Inc, 2002.
- Von Hippel, P.T. “Achievement, learning, and seasonal impact as measures of school effectiveness: it’s better to be valid than reliable,” School Effectiveness and School Improvement. 20(2), 187-213, 2009.
- Von Secker. C. E., & Lissitz. R. W. “Estimating school value-added effectiveness: consequence of respecification of hierarchical linear model,” Paper presented at the American Educational Research Association Annual Meeting. Chicago, 1997.