

การประยุกต์เทคนิคการวิเคราะห์การสมนัย ในการวิจัยทางสังคมศาสตร์ Correspondence Analysis Applied to Social Science Research

ประสพชัย พสุนนท์ (Prasopchai Pasunon)¹

อาฟีฟี ลาเต๊ะ (Afifi Lateh)²

เกตุวดี สมบูรณ์ทวี (Kedwadee Sombultawee)^{3*}

บทคัดย่อ

บทความนี้มีวัตถุประสงค์ในการนำเสนอการวิเคราะห์การสมนัย (Correspondence Analysis: CA) เพื่อใช้ในการวิจัยทางสังคมศาสตร์ โดยเน้นหนักไปที่หลักการ แนวคิด วิธีการคำนวณ และโปรแกรมสำเร็จรูปที่ช่วยในการวิเคราะห์ เนื่องจากวิธีการ CA ถูกมองว่าเป็นสถิติขั้นสูงและยุ่งยากต่อการนำไปใช้ ดังนั้น บทความนี้จึงได้พยายามทำความเข้าใจถึงที่มาและวิธีการในการนำ CA ไปประยุกต์ในงานทางสังคมศาสตร์ โดยเฉพาะการตีความข้อมูลในระดับนามบัญญัติ หรือตัวแปรเชิงกลุ่ม ซึ่งผู้วิจัยโดยทั่วไปนิยมใช้สถิติการทดสอบไคสแควร์ ในการทดสอบความสัมพันธ์ระหว่างตัวแปร อย่างไรก็ตาม หากผู้วิจัยนำวิธีการ CA เข้ามาร่วมในการวิเคราะห์ด้วยแล้วจะทำให้ได้สารสนเทศที่น่าสนใจจากผลการวิจัยมากขึ้น

คำสำคัญ : การวิเคราะห์ความสมนัย การวิจัยทางสังคมศาสตร์

¹รองศาสตราจารย์ประจำคณะวิทยาการจัดการ มหาวิทยาลัยศิลปากร วิทยาเขตสารสนเทศเพชรบุรี

²ผู้ช่วยศาสตราจารย์ประจำคณะศึกษาศาสตร์ มหาวิทยาลัยสงขลานครินทร์ วิทยาเขตปัตตานี

³อาจารย์ประจำคณะวิทยาการจัดการ มหาวิทยาลัยศิลปากร วิทยาเขตสารสนเทศเพชรบุรี

*Corresponding author, E-mail : kedwadee_so@hotmail.com

Abstract

This paper aimed to demonstrate that correspondence analysis for social science research by focusing on the principles, concepts, computational method and how to perform CA using SPSS software. As correspondence analysis was seen as the advance statistic and complexity to use. Thus, this article provided the understanding of background and the method for correspondence analysis in order to apply for social science research. Especially for the nominal scale data interpretation or categorical variables which general researchers typically used Chi-square statistics technique to examine the relationship between variables. However, correspondence analysis can be used as supplementary technique in the data analysis and to make the result more interesting and robustness.

Keywords : Correspondence analysis, Social science

บทนำ

การวิจัยทางสังคมศาสตร์ (Social Science Research) ที่ใช้ระเบียบวิธีวิทยาการวิจัยเชิงปริมาณ (Qualitative Research Methodology) ที่มีการเก็บข้อมูลเชิงกลุ่ม (Categorical Data) หรือข้อมูลมาตราชานามบัญญัติ (Nominal Scale) ดังเช่น เพศ อาชีพ ตราสินค้า ภูมิลำเนา มักจะแสดงเป็นลักษณะทางประชากรศาสตร์ เพื่ออธิบายลักษณะของกลุ่มตัวอย่างที่ได้สำรวจ หรือเก็บข้อมูลมา หรืออาจจะนำเสนอด้วยการวิเคราะห์ข้อมูลที่อยู่ในรูปตารางการถ่วงสองทาง (Two-way Contingency Table) เพื่อทดสอบความสัมพันธ์ระหว่างตัวแปรทางแถวกับตัวแปรทางคอลัมน์ว่ามีนัยสำคัญทางสถิติหรือไม่ ซึ่งการวิเคราะห์ข้อมูลดังกล่าวจะไม่สามารถอธิบายความสัมพันธ์ระหว่างระดับต่างๆ ของตัวแปรสองกลุ่มข้างต้น รวมทั้งไม่สามารถอธิบายความสัมพันธ์ระหว่างระดับภายในตัวแปรเดียวกันได้ (Beh, 2004; Doey and Kurta, 2011)

การวิเคราะห์การสมนัย (Correspondence Analysis: CA) เป็นวิธีการวิเคราะห์ทางสถิติเชิงพหุในการค้นหาหรืออธิบายข้อมูลเชิงความถี่ของตัวแปรเชิงกลุ่ม (Categorical Variables) ในรูปตารางการถ้อยสองทางหรือมากกว่า ซึ่งจะสามารถอธิบายความสัมพันธ์ระหว่างตัวแปรทั้งในระดับภาพรวม (Overall Association) ความสัมพันธ์ระหว่างระดับของตัวแปร (Association of Categories Between Variable) และความสัมพันธ์ระหว่างระดับภายในตัวแปร (Association of Categories Within a Variable) อีกทั้งยังช่วยในการสร้างแผนภาพความคิดจากข้อมูล (Data Visualization) (ชยุตม์ ภิรมย์สมบัติ, 2557)

โดยทั่วไป การวิเคราะห์การสมนัยแบ่งออกเป็น 2 ประเภท คือ 1) การวิเคราะห์การสมนัยอย่างง่าย (Simple Correspondence Analysis) เป็นการวิเคราะห์ตัวแปรเชิงกลุ่ม 2 ตัวแปร เช่น ความสัมพันธ์ระหว่างมหาวิทยาลัย กับภาพลักษณ์ของมหาวิทยาลัย (นลินี พานสายตา, 2555) หรือรูปแบบการเรียนรู้ กับลักษณะแรงจูงใจ (Askeff-Williams and Lawson, 2004) หรือความสัมพันธ์ระหว่างกลุ่มคำกับบริบทภาษาศาสตร์ (Bertin and Atanassova, 2015) หรือความสัมพันธ์ระหว่างเพศและกลุ่มอายุกับประเภทของภาพยนตร์ (Redfern, 2012) และ 2) การวิเคราะห์การสมนัยพหุคูณ (Multiple Correspondence Analysis) เป็นการวิเคราะห์ตัวแปรเชิงกลุ่มตั้งแต่ 3 ตัวแปรขึ้นไป เช่น ความสัมพันธ์ระหว่างกลุ่มอายุ สถานภาพ ระดับการศึกษา อาชีพ ความตั้งใจที่จะไปทำบุญ กับระดับความคิดเห็นต่อการกำหนดให้วันอาทิตย์เป็นวันทำบุญ (วิยะดา ต้นวัฒนากุล และคณะ, 2548) หรือความสัมพันธ์ระหว่างประเภทรถยนต์ สภาพรถยนต์ วิธีการซื้อ แหล่งข้อมูล การตรวจสอบสภาพก่อนซื้อ และการโอนกรรมสิทธิ์ (เกษิตศ แสงวิรุณและกฤษฎ จรินทร์, 2555) เป็นต้น

เนื่องจากวิธีการ CA ถูกมองว่าเป็นสถิติขั้นสูงและยุ่งยากต่อการนำไปใช้ ดังนั้น บทความนี้มีวัตถุประสงค์เพื่อนำเสนอและอธิบายหลักการและแนวคิดของการวิเคราะห์การสมนัย จากตัวอย่างงานวิจัยด้วยการวิเคราะห์การสมนัย ตลอดจนการประยุกต์ด้วยโปรแกรมคอมพิวเตอร์นำไปสู่การอธิบายความสัมพันธ์ของตัวแปรเชิงกลุ่มที่ให้อารมณ์เห็นได้อย่างเจาะลึกในแต่ละระดับของความสัมพันธ์ของงานวิจัย

ทางสังคมศาสตร์ อีกทั้งช่วยให้นักวิจัยสามารถนำเทคนิค CA ไปประยุกต์ใช้ในงานวิจัยเพื่อให้ได้สารสนเทศที่น่าสนใจจากผลการวิจัยมากขึ้น

หลักการและแนวคิดของการวิเคราะห์การสมนัย

วิธีการ CA ได้รับการพัฒนาขึ้นเมื่อกว่า 50 ปีที่ผ่านมา โดยมีชื่อเรียกที่หลากหลาย อาทิ Dual Scaling, Reciprocal Averages, Categorical Discriminant Analysis Method of Reciprocal Averages, Guttman Weighting, Principal Component Analysis of Qualitative Data และ Biplot Analysis (Richardson M., Kuder G. F. 1933; Horst, 1935; Guttman, 1953; Hoffman and Franke, 1986) โดยชื่อเหล่านี้เกิดจากวัตถุประสงค์หรือมีคำถามการวิจัยที่แตกต่างกันสำหรับผู้ให้กำเนิดวิธีการ CA อย่างเห็นทาง การคือ Jean-Paul Benzecri นักสถิติชาวฝรั่งเศส โดยดำรงตำแหน่งศาสตราจารย์ ณ มหาวิทยาลัย Pierre-et-Marie-Curie ในกรุงปารีส (Armatté, 2008) โดยในปี ค.ศ. 1973 ศาสตราจารย์ Benzecri ได้เขียนบทความชื่อ 'l' analyse factorielle des correpondeances ก่อนที่งานดังกล่าวจะได้ถูกเรียบเรียงเป็นภาษาอังกฤษโดย Hill ในปี ค.ศ. 1974 (Hill and Gauch, 1980)

วิธีการ CA ถูกนำเสนอในงานวิจัยทางการตลาดในปี 1960 และปี 1970 ก่อนจะได้รับแพร่หลายในสาขาวิชาอื่นๆ ภายหลังที่เผยแพร่เป็นภาษาอังกฤษ ทำ CA ถูกนำไปประยุกต์ใช้และเป็นที่ยอมรับในทางวิชาการมากขึ้น (Clausen, 1988) ปัจจุบันวิธีการ CA มีการนำไปประยุกต์ใช้ในหลายสาขา เช่น เศรษฐศาสตร์ เคมี การศึกษา การเงิน ฯลฯ ราชบัณฑิตยสถาน (2558) บัญญัติภาษาไทยของ Correspondence ในสาขาปรัชญาไว้ว่า “ความสมนัย” (2 มีนาคม 2545) และในสาขาคณิตศาสตร์ไว้ว่า “การสมนัย” (17 กรกฎาคม 2547) โดย Correspondence เป็นคำนามของ Correspond จาก รากศัพท์ respondere “to respond, to answer” รวมกับ Cor- “together, into or in agreement” แปลความได้ว่า การตอบสนองที่มีความเหมือนกันหรือสอดคล้องกัน (ชยุตม์ ภิรย์สมบัติ, 2557)

ดังได้กล่าวในตอนต้นว่าการวิเคราะห์การสมนัยเป็นการค้นหาหรืออธิบาย ข้อมูลเชิงความถี่ของตัวแปรเชิงกลุ่ม (Categorical Variables) ในรูปตารางการถัวจร สองทางหรือมากกว่า กล่าวคือ เป็นการแปลงความถี่ (Frequency) ในตารางการถัว จร (Contingency) ให้เป็นระยะทางเชิงสถิติ (Statistical Distance) ที่สามารถให้ ความสอดคล้องกันระหว่างกันของแถว (Row) และคอลัมน์ (Column)ในระดับ ย่อยผ่านระยะทางดังกล่าว โดยทั่วไปนิยมใช้ระยะทางไคสแควร์ (Chi-square Distance) ซึ่งจะนำเสนอในรูปแบบภาพเพื่อแสดงความสัมพันธ์ของตัวแปรทางแถว และตัวแปรทางคอลัมน์ รวมทั้งความสัมพันธ์ระหว่างระดับของตัวแปรทางแถวและ ทางคอลัมน์ และความสัมพันธ์ระหว่างระดับภายในตัวแปรทางแถวและทางคอลัมน์ ด้วยแผนภาพการสมนัย (Correspondence Map) (Hair et al., 2010) หรืออาจ เรียกว่าแผนภาพเชิงการรับรู้ (Perceptual Map) หรือแผนภาพทวิ (Biplot) ส่งผลให้ ผู้วิจัยสามารถตีความความสัมพันธ์ได้ง่ายผ่านระยะทางระหว่างจุด กล่าวคือ หากจุด (ซึ่งเป็นตัวแทนระดับย่อย) ที่อยู่ใกล้กัน แสดงว่ามีความสอดคล้องหรือ สัมพันธ์กันมากกว่าจุดที่อยู่ห่างกัน นอกจากนี้ CA ยังช่วยในการลดมิติ (Dimension Reduction) ของข้อมูลจากเดิม โดยไม่มีการทดสอบนัยสำคัญทางสถิติ

สำหรับข้อตกลงเบื้องต้นของการวิเคราะห์การสมนัยนั้นต้องเข้ากับข้อมูล ที่ไม่ต่อเนื่อง (Discrete Data) หรือหากข้อมูลเดิมเป็นข้อมูลต่อเนื่อง (Continuous Data) ต้องทำให้เป็นข้อมูลไม่ต่อเนื่องก่อน รวมทั้งความถี่ในแต่ละช่องของตาราง การถัวจรต้องไม่มีค่าเป็นลบ (Non-negative Values) และควรใช้กับตัวแปรเชิง กลุ่มที่มีจำนวนหลายกลุ่ม ในกรณีที่มีเพียง 2 หรือ 3 กลุ่ม ผลการวิเคราะห์ด้วย CA จะให้ผลลัพธ์ของสารสนเทศที่ไม่แน่นอน หรือให้สารสนเทศที่ไม่ดีไปกว่าตาราง การถัวจรเดิมการวิเคราะห์ด้วย CA ไม่ได้มีข้อตกลง (Assumption) เกี่ยวกับการ แจกแจงความน่าจะเป็น (Probability Distribution) โดยเฉพาะว่าต้องมีการแจกแจง แบบปกติ (Normal Distribution) แต่หากในตารางการถัวจรมีค่าผิดปกติ (Outliers) ปะปนในข้อมูลจะส่งผลที่ได้จาก CA เบี่ยงเบนไปจากที่ควรจะเป็น (Bendixen, 2003) นอกจากนี้ CA ควรมีคุณสมบัติความคล้ายคลึงกันระหว่างแถวและคอลัมน์ เช่น ในแต่ละช่อง (Cell) ต้องไม่มีช่องว่าง เป็นต้น (Doey and Kurta, 2011)

ข้อมูลในการวิเคราะห์การสมนัยแสดงในรูปตารางการถัวสองทาง (Two-way Contingency Table) เรียกว่า ตารางการสมนัย (Correspondence Table) ซึ่งเป็นเมทริกซ์ (Matrix) ของข้อมูลเดิมที่นำเสนอการแจกแจงความถี่ n_{ij} บางครั้งเรียกตาราง Primitive หรือเมทริกซ์ Primitive (กมล สนิทธรรม, 2549) (เมทริกซ์ Primitive เป็นเมทริกซ์ที่ไม่มีค่าเป็นลบ) โดยในโปรแกรม SPSS นำเสนอผลรวมของแถว ($n_{.j}$) และผลรวมของคอลัมน์ ($n_{i.}$) ในชื่อ Active Margin โดยที่ n_{ij} คือ ข้อมูลจำนวนนับหรือความถี่ที่สังเกตได้ของตัวแปรการวิจัย ในแถวที่ i และคอลัมน์ที่ j ($i = 1, 2, \dots, r; j = 1, 2, \dots, c$) เมื่อ $n_{.j} = \sum_{i=1}^r n_{ij}$ $n_{i.} = \sum_{j=1}^c n_{ij}$ และ $n = \sum_{i=1}^r \sum_{j=1}^c n_{ij}$ แสดงดังตารางที่ 1

การวิเคราะห์การสมนัยแสดงด้วยจำนวนมิติ (Dimension) เพื่อการอธิบายปรากฏการณ์ที่เกิดขึ้นจากตัวแปรเชิงกลุ่ม โดยจำนวนมิติสูงสุดสามารถพิจารณาได้จาก $\min(r-1, c-1)$ เมื่อ $\min$ คือ ค่าต่ำสุด สำหรับจำนวนมิติขึ้นกับร้อยละของความเฉื่อยทั้งหมด (% of Total Inertia) โดยทั่วไปนิยมที่ประมาณ 80% หรือ 90% โดยแต่ละมิติควรมีความเฉื่อยมากกว่า $\frac{1}{\min(r,c)-1} \%$ ของความเฉื่อยรวม หรืออาจพิจารณาจาก Scree Plot ของความเฉื่อยที่มีอัตราการลดลงอย่างเด่นชัด

สำหรับความเฉื่อย (Inertia) หรือค่าไอเกน (Eigenvalues) ซึ่งมีความหมายในลักษณะเดียวกันกับ “โมเมนต์ความเฉื่อย” (Moment of Inertia) ตามหลักคณิตศาสตร์ประยุกต์ (พบในสาขาวิศวกรรมศาสตร์) กล่าวคือ เป็นการหาปริพันธ์ (Integral) ระยะทางกำลังจากมวล (Mass) ถึง Centroid (Greenacre, 1984) ความเฉื่อยจึงเป็นผลรวมของเพียร์สันไคสแควร์จากตารางการถัวสองทางที่สะท้อนถึงความสัมพันธ์ในแต่ละมิติที่มีระดับความสำคัญจากมากไปน้อย โดยสามารถคำนวณความเฉื่อยได้ดัง (1) และความเฉื่อยรวม (Total Inertia) คำนวณจาก (2) ซึ่งเป็นผลรวมของค่าไอเกนที่ได้จากการกระจายของจุด (หรือมวล) รอบๆ Centroid เมื่อ m_i แทนมวลของจุดที่ i และ d_i^2 แทนระยะทางแบบเพียร์สันไคสแควร์จากจุดที่ i ถึง Centroid โดยความเฉื่อยรวมมีความสัมพันธ์ใกล้ชิดกับสถิติไคสแควร์ (นั่นคือ $\text{Total Inertia} = \frac{\chi^2}{n}$) เมื่อ $P_{ij} = \frac{n_{ij}}{n}$

$$\text{Inertia} = \sum_{i=1}^c \sum_{j=1}^r \frac{(p_{ij} - p_{i\cdot} \cdot p_{\cdot j})^2}{p_{i\cdot} \cdot p_{\cdot j}} = md^2 \quad \text{_____}(1)$$

$$\text{Total Inertia} = md_i^2 \quad \text{_____}(2)$$

สำหรับค่าเอกฐาน (Singular Values) เป็นการอธิบายความสัมพันธ์เชิงคาโนนิคอลสูงสุดระหว่างระดับของตัวแปรในการวิเคราะห์มิติแต่ละมิติ ซึ่งสามารถคำนวณได้จากค่ารากที่สองของความเฉื่อยในแต่ละมิติ และสัดส่วนความเฉื่อย (Proportion of Inertia) คำนวณได้จากค่าสัดส่วนของความเฉื่อยแต่ละมิติต่อความเฉื่อยรวม

การทดสอบสมมติฐานของการวิเคราะห์ความสมนัย

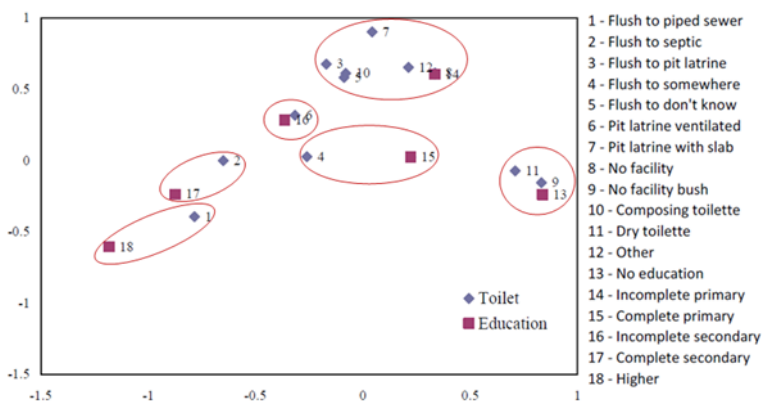
การทดสอบสมมติฐานของ CA โดยทั่วไปจะนิยมพิจารณาความเป็นอิสระกันของตัวแปรทางแถวกับตัวแปรทางคอลัมน์สามารถกำหนดสมมติฐานการทดสอบได้จากสมมติฐานหลัก (H_0) คือ ตัวแปรทางแถวกับตัวแปรทางคอลัมน์เป็นอิสระกัน และสมมติฐานแย้ง (H_1) คือ ตัวแปรทางแถวกับตัวแปรทางคอลัมน์ไม่เป็นอิสระกัน โดยใช้สถิติทดสอบไคสแควร์ (Chi-square Statistic: χ^2) ดัง (3) มีการแจกแจงแบบไคสแควร์ (Chi-square Distribution) โดยที่ $E_{ij} = \frac{n_{i\cdot} \cdot n_{\cdot j}}{n}$ และองศาความเป็นอิสระ (Degree of freedom: df) เท่ากับ $(r - 1)(c - 1)$

$$\chi^2 = \chi^2 = \sum_{i=1}^c \sum_{j=1}^r \frac{(n_{ij} - E_{ij})^2}{E_{ij}} \quad \text{_____}(3)$$

นอกจากนี้ CA ยังสามารถอธิบายน้ำหนักความสำคัญของระดับของตัวแปรในแต่ละมิติหรือการตีความความสัมพันธ์ที่ได้ว่าสามารถอธิบายระดับของตัวแปรนั้นๆ ได้ดีเพียงใดอีกด้วย สำหรับการอธิบายความสัมพันธ์ระหว่างตัวแปรในระดับภาพรวมสามารถอธิบายได้จากสถิติทดสอบไคสแควร์ แต่การอธิบายความสัมพันธ์ระหว่างระดับของตัวแปร และความสัมพันธ์ระหว่างระดับภายในตัวแปร สามารถพิจารณาได้จากแผนภาพทวิที่แสดงการรับรู้ที่พลอตจากค่าของคะแนนในแต่ละมิติซึ่งเป็นการอธิบายความสัมพันธ์โดยอาศัยตำแหน่งความใกล้ชิด นั่นคือ ระดับของตัวแปรใดๆ

ที่มีตำแหน่งความสัมพันธ์ใกล้เคียงกันมากแสดงว่าระดับตัวแปรนั้นๆ มีความสัมพันธ์กันสูง บางครั้งการตีความด้วยแผนภาพการรับรู้อาจจะต้องอาศัยประสบการณ์ในงานวิจัยนั้นๆ ร่วมด้วย

เพื่อให้เข้าใจได้ง่ายขึ้น ผู้เขียนใช้รูปที่ 1 ประกอบในการอธิบายโดยเป็นการศึกษากลุ่มของระดับการศึกษา กับ พฤติกรรมการใช้ห้องน้ำในประเทศอินเดียของ Wei Pillai and Maleku (2014) ซึ่งสามารถแบ่งกลุ่มได้เป็น 6 กลุ่ม คือ กลุ่ม 1 เป็นกลุ่มที่ไม่ได้เรียนหนังสือจะถ่ายตามพุ่มไม้ หรือห้องน้ำแบบแห้ง กลุ่ม 2 เป็นกลุ่มเรียนไม่จบระดับประถมศึกษาจะขุดหลุมถ่าย ใช้ส้วมหลุม หรืออื่นๆ กลุ่ม 3 เป็นกลุ่มเรียนจบระดับประถมศึกษาจะถ่ายที่ใดที่หนึ่ง กลุ่ม 4 เป็นกลุ่ม เรียนไม่จบระดับมัธยมศึกษา จะใช้ส้วมหลุมระบายอากาศ กลุ่ม 5 เป็นกลุ่มเรียนจบระดับมัธยมศึกษาจะใช้ส้วมแบบชักโครก กลุ่ม 6 เป็นกลุ่มที่จบสูงกว่าระดับมัธยมศึกษาจะใช้ส้วมแบบชักโครกแบบมีการระบายลงท่อ เห็นได้ว่าเมื่อพิจารณาจากความใกล้ชิดของความสัมพันธ์แล้วอาจจะสามารถรวมกลุ่ม 3 และ 4 หรือกลุ่ม 5 และ 6 เข้าด้วยกันได้



รูปที่ 1 การศึกษากลุ่มของระดับการศึกษา กับ พฤติกรรมการใช้ห้องน้ำในประเทศอินเดีย

ที่มา : ปรับจาก Wei Pillai and Maleku (2014)

จากแผนภาพการรับรู้ในรูปที่ 1 เห็นได้ว่าผู้วิจัยสามารถพิจารณาความสัมพันธ์ระหว่างระดับของตัวแปร รวมทั้งความสัมพันธ์ระหว่างระดับภายในตัวแปรที่อธิบายสารสนเทศเชิงลึกเพิ่มมากขึ้นเมื่อเทียบกับการทดสอบด้วยสถิติไคสแควร์จากตารางการถ้อยจรรยาบรรณในรูปแบบเดิมที่เคยวิเคราะห์และนำเสนอ และแผนภาพการรับรู้ข้างต้นยังสามารถรวมระดับของตัวแปรที่มีลักษณะใกล้เคียงหรืออาจเรียกว่าสามารถลดจำนวนตัวแปรลงเพื่อทำให้การวิเคราะห์ทางสังคมศาสตร์ด้วยสถิติอื่นในขั้นตอนต่อไปง่ายขึ้น และให้สารสนเทศที่ถูกต้อง ชัดเจนมากขึ้นด้วย

การวิเคราะห์การสมนัยด้วยโปรแกรมคอมพิวเตอร์

บทความนี้ เลือกใช้ข้อมูลตัวอย่างตัวแปรเชิงกลุ่ม 2 ตัวแปรของ Bendixen (2003) เพื่อทำความเข้าใจในวิธีการในเบื้องต้นของ CA โดยเป็นข้อมูลเกี่ยวข้องกับจำนวนผู้ใช้และการเลือกใช้ยาสีฟัน 4 ตราสินค้า (Brand) คือ A B C และ D ใน 3 พื้นที่ (Region) คือ พื้นที่ที่ 1 – 3 ดังตารางที่ 4 โดยมีขนาดตัวอย่าง (n) เท่ากับ 120 เพื่อนำมาวิเคราะห์ด้วยโปรแกรมสำเร็จรูปทางสถิติ SPSS

สำหรับขั้นตอนการวิเคราะห์ CA ด้วยโปรแกรม SPSS มีดังนี้

1. นำข้อมูลตารางการถ้อยจรรยาบรรณจากตารางที่ 4 ลงไปใน Data View ของ SPSS โดยระดับ 1 – 4 ของ Brand คือ A B C และ D ตามลำดับ และระดับ 1 – 3 ของ Region คือ พื้นที่ที่ 1 – 3 ตามลำดับ ดังรูปที่ 2 ในกรณีนี้เป็นการนำข้อมูลแบบความถี่ จึงต้องถ่วงน้ำหนักก่อน โดยเลือกเมนู Data >>> Weight Cases จากนั้นคลิก Weight cases by แล้วเลือก Frequency เข้าไปในช่อง Frequency Variable: ดังรูปที่ 3 จากนั้นคลิก OK อย่างไรก็ตาม หากเป็นการนำเข้าแบบทั่วไปไม่จำเป็นต้องทำการถ่วงน้ำหนัก

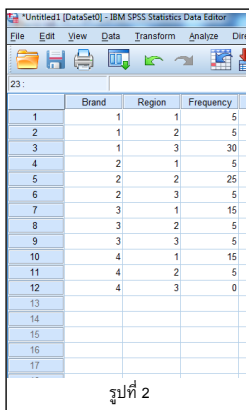
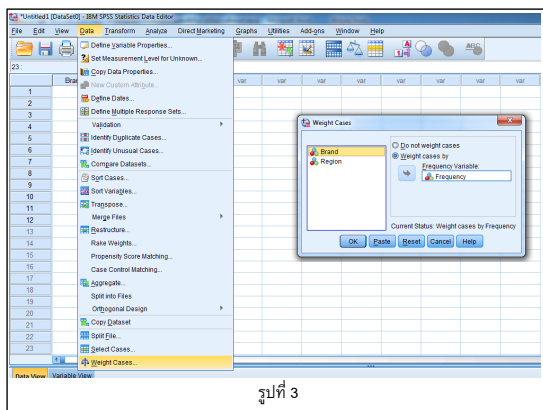


Figure 2 shows a dataset with the following data:

	Brand	Region	Frequency
1	1	1	5
2	1	2	5
3	1	3	30
4	2	1	5
5	2	2	25
6	2	3	5
7	3	1	15
8	3	2	5
9	3	3	5
10	4	1	15
11	4	2	5
12	4	3	0
13			
14			
15			
16			
17			

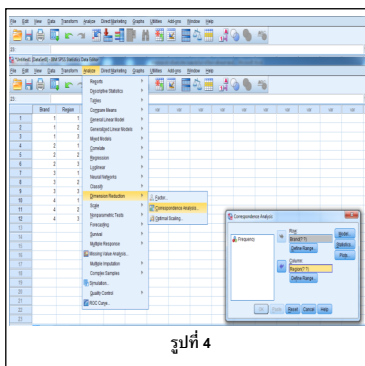
รูปที่ 2



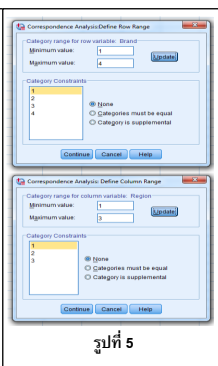
รูปที่ 3

2. เลือกเมนู Analyze >>> Dimension Reduction >>>

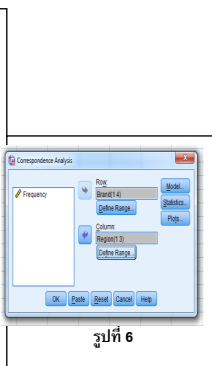
Correspondence Analysis ปรากฏ Dialog Box ให้นำตัวแปร Brand เข้าในช่อง Row: และนำตัวแปร Region เข้าในช่อง Column: ดังรูปที่ 4 จากนั้นให้ Define Range โดยในส่วนของ Row: ให้ใส่ 1 ในช่อง Minimum value และใส่ 4 ในช่อง Maximum value คลิก Update และคลิก Continue ส่วน Column : ทำเช่นเดียวกับ Row: แต่ให้ใส่ 1 ในช่อง Minimum value และใส่ 3 ในช่อง Maximum value ดังรูปที่ 5 เมื่อดำเนินการแล้วจะได้ดังรูปที่ 6



รูปที่ 4



รูปที่ 5



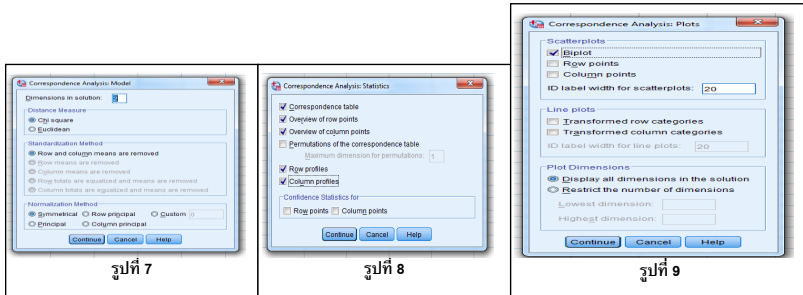
รูปที่ 6

สำหรับรูปที่ 5 มีตัวเลือกในการกำหนดเงื่อนไขของแต่ละระดับ 3 แบบ คือ 1) None คือ ไม่กำหนดเงื่อนไขใดๆ 2) Categories must be equal เป็นการให้เงื่อนไขระดับย่อยมีน้ำหนักในการวิเคราะห์ CA เท่ากัน และ 3) Categories is supplemental เป็นการให้เงื่อนไขระดับย่อยบางระดับไม่ต้องมีผลต่อ CA แต่ยังมีผลต่อแผนภาพการสมนัย นั่นคือ supplemental จะไม่มีผลต่อมิติที่ได้จาก CA

3. จากรูปที่ 6 ผู้วิจัยสามารถกำหนดรายละเอียดการวิเคราะห์ของ CA โดยเลือกคลิก Model หรือ Statistics หรือ Plots ดังนี้

3.1 รูปที่ 7 เป็นผลจากการคลิก Model ในรูปที่ 5 โดยในช่อง Dimensions in solution: เป็นการกำหนดมิติจาก CA โดยทั่วไปกำหนดที่ 2 หรือ 3 มิติ เพราะง่ายต่อการตีความ (ตัวอย่างนี้วิเคราะห์ 2 มิติ) ในส่วน Distance Measure สามารถคลิกเลือก Chi square หรือ Euclidian (โดยทั่วไปนิยมใช้ระยะทางไคสแควร์) และสำหรับ Standardization Method ให้คลิกเลือก Row and column means are removed เพื่อให้ CA เป็นแบบมาตรฐานที่มีค่าเฉลี่ยของแถวและคอลัมน์ถูกลบออกจากแต่ละแถวและแต่ละคอลัมน์ โดยยังสามารถเลือกคลิกการลบเฉพาะค่าเฉลี่ยของแถวหรือของคอลัมน์ และการปรับความถี่รวมตามแถวหรือตามคอลัมน์ก่อนลบค่าเฉลี่ยออก

สำหรับ Normalization Method นิยมเลือกคลิก Symmetrical ซึ่งให้ผล CA ระหว่างระดับย่อยๆ ของแถวและคอลัมน์พร้อมๆ กัน หากเลือก Principal คือ ต้องการผล CA ระดับย่อยของแถวและคอลัมน์ แต่ถ้าเลือก Row Principal หรือ Column Principal คือ ต้องการผล CA ระดับย่อยภายในแถวหรือภายในคอลัมน์ตามลำดับ ในส่วนของ Custom ต้องระบุตัวเลขตั้งแต่ -1 ถึง 1 หากกำหนด -1, 0 และ 1 จะให้ผลลัพธ์เหมือนกับ Column Principal, Symmetrical และ Row Principal ตามลำดับ หากกำหนดเป็นค่าอื่นๆ จะให้ค่าการกระจายความแปรปรวนที่ใช้อธิบายตัวแปรแถวและคอลัมน์ที่แตกต่างกันไป



3.2 รูปที่ 8 เป็นผลจากการคลิก Statistics ในรูปที่ 6 โดยถ้าคลิก Correspondence table จะให้ตารางการถ่วงน้ำหนักของข้อมูล หากคลิก Overview of row points จะแสดงตารางสรุปของตัวแปรแถวที่ประกอบด้วยคะแนนของแต่ละระดับในแต่ละมิติ (Score in Dimension) ค่าน้ำหนัก (Mass) ค่าความแปรปรวนที่อธิบายได้ (Inertia) ของแต่ละระดับ และส่วนของความแปรปรวนที่อธิบายได้ในแต่ละมิติ (Contribution of Point to Inertia of Dimension) แต่หากคลิก Overview of column points ให้ผลลัพธ์เช่นเดียวกันแต่เปลี่ยนเป็นตัวแปรคอลัมน์ สำหรับการคลิก Row Profiles และ Column Profiles จะแสดงตารางการถ่วงน้ำหนักของข้อมูลเฉลี่ยตามแถวและคอลัมน์ ตามลำดับ

ผู้วิจัยยังเลือกคลิก Permutations of the Correspondence table เพื่อแสดงตารางการถ่วงน้ำหนักโดยจัดเรียงแถวและคอลัมน์จากคะแนนน้อยไปมาก โดยสามารถระบุได้ว่าต้องการเรียงคะแนนจำนวนกี่มิติ หากไม่ระบุจะแสดงเพียงมิติแรกที่วิเคราะห์ได้ นอกจากนี้ การคลิก Confidence Statistics for Row points จะให้ช่วงความเชื่อมั่นของ CA พร้อมทั้งค่าสัมประสิทธิ์และส่วนเบี่ยงเบนมาตรฐานของตัวแปรแถว แต่หากคลิก Confidence Statistics for Column points ก็จะให้ผลลัพธ์เช่นเดียวกันเพียงแค่เป็นตัวแปรคอลัมน์

3.3 รูปที่ 9 เป็นผลจากการคลิก Plots ในรูปที่ 6 โดยในส่วนของ Scatterplots การคลิก Biplot จะแสดงผลการวิเคราะห์ในแผนภาพเดียวกัน แต่หากคลิก Row points และ Column point จะแสดงแผนภาพแยกตามแถวและคอลัมน์ตามลำดับ การคลิก Line plots เป็นการแสดงค่าของระดับจากข้อมูลดั้งเดิมเปรียบเทียบกับค่าของระดับที่ปรับรูป (Transformed) แล้วตามแถวหรือคอลัมน์ ในส่วน Plot Dimension นิยมคลิก Display all dimension in the solution

การแปลผลการวิเคราะห์ข้อมูล

เมื่อดำเนินการวิเคราะห์ข้อมูลด้วย CA ตามขั้นตอนตามรูปที่ 2 – 9 แล้วจะได้ผลลัพธ์ตามรูปที่ 10– 14 ซึ่งผลลัพธ์จาก SPSS โดยสามารถอธิบายได้ดังนี้

1. รูปที่ 10 คือ ตารางการกระจายระหว่างตราสินค้าและพื้นที่ โดย CA จะใช้คำว่า Active Margin แทน Total เพราะในโปรแกรมสามารถแถวหรือคอลัมน์ให้เป็น supplemental ได้ ในส่วนของรูปที่ 11 คือ ตารางสัดส่วนตามแถวและตารางสัดส่วนตามคอลัมน์ เป็นการแสดงค่าน้ำหนักของแถวและคอลัมน์ โดยจะเห็นว่าพื้นที่ 1 ให้น้ำหนักกับตราสินค้า C และ D ในขณะที่พื้นที่ 2 ให้น้ำหนักกับตราสินค้า B ส่วนพื้นที่ 3 เน้นให้ความสำคัญกับตราสินค้า A

Brand	Region			
	Region 1	Region 2	Region 3	Active Margin
Brand A	5	5	30	40
Brand B	5	25	5	35
Brand C	15	5	5	25
Brand D	15	5	0	20
Active Margin	40	40	40	120

รูปที่ 10

Brand	Region			
	Region 1	Region 2	Region 3	Active Margin
Brand A	.125	.125	.750	1.000
Brand B	.143	.714	.143	1.000
Brand C	.600	.200	.200	1.000
Brand D	.750	.250	.000	1.000
Mass	.333	.333	.333	

Brand	Region			
	Region 1	Region 2	Region 3	Mass
Brand A	.125	.125	.750	.333
Brand B	.125	.625	.125	.292
Brand C	.375	.125	.125	.208
Brand D	.375	.125	.000	.167
Active Margin	1.000	1.000	1.000	

รูปที่ 11

2. จากรูปที่ 12 จะพบว่า $\chi^2 = 79.607$ ด้วยองศาความเป็นอิสระ (Degree of Freedom: df) เท่ากับ 6 มีค่า Sig. (หรือ P-value) เท่ากับ 0.000 นั่นคือ ตราสินค้าและพื้นที่มีความสัมพันธ์อย่างมีนัยสำคัญ 0.01 และพบมิติจาก CA จำนวน 2 มิติ มีค่า Total Inertia = 0.663 (คำนวณจาก $\chi^2/n = 79.607/120$) หมายความว่า เมื่อรวมทุกมิติเข้าด้วยกันสามารถอธิบายความแปรปรวนได้ร้อยละ 66.3 สำหรับมิติที่ 1 มีค่าความเฉื่อยเท่ากับ 0.41 โดย $\sqrt{0.4} = 0.641 =$ Singular Value (ในมิติอื่นๆ ก็มีการคำนวณเช่นเดียวกัน) โดยทั่วไปนิยมตีความค่าความแปรปรวนที่ใช้อธิบายในแต่ละมิติจากสัดส่วนของความเฉื่อย ซึ่งพบว่า มิติที่ 1

และ 2 อธิบายความแปรผันได้ร้อยละ 61.8 และ 38.2 ของความผันแปรรวมร้อยละ 66.3 ตามการวิเคราะห์การสมนัย ตามลำดับ (โดยที่การคำนวณในมิติที่ 2 คือ $0.382 = 0.253/0.663 = \text{Inertia}/(\text{Total Inertia})$ สำหรับการคำนวณในมิติอื่นๆ ก็เช่นกัน) สำหรับค่า Standard Deviation เป็นการพิจารณาถึงความแม่นยำของความสัมพันธ์ที่แสดงด้วย Singular Value ในแต่ละมิติ

Summary								
Dimension	Singular Value	Inertia	Chi Square	Sig.	Proportion of Inertia		Confidence Singular Value	
					Accounted for	Cumulative	Standard Deviation	Correlation 2
1	.641	.410			.618	.618	.069	.180
2	.503	.253			.382	1.000	.086	
Total		.663	79.607	.000 ^a	1.000	1.000		

a. 6 degrees of freedom

รูปที่ 12

3. รูปที่ 13 เป็นตารางภาพรวมของค่าน้ำหนัก (Mass) จากรูปที่ 10 แสดงคะแนนในแต่ละมิติเพื่อนำไปสร้างแผนภาพการวิเคราะห์การสมนัย โดยแสดงค่าความเฉื่อยในแต่ละมิติ พร้อมทั้งแสดงสัดส่วนของความแปรผันในมิติที่ 1 และ 2 ที่ใช้อธิบายด้วยระดับของตัวแปร และสัดส่วนของความแปรผันในแต่ละระดับที่อธิบายด้วยมิติที่ 1 และ 2

Overview Row Points ^a									
Brand	Mass	Score in Dimension		Inertia	Contribution				
		1	2		Of Point to Inertia of Dimension		Of Dimension to Inertia of Point		Total
					1	2	1	2	
Brand A	.333	1.097	.148	.260	.626	.015	.986	.014	1.000
Brand B	.292	-.397	-1.047	.190	.072	.636	.155	.845	1.000
Brand C	.208	-.424	.638	.067	.058	.169	.359	.641	1.000
Brand D	.167	-.968	.739	.146	.244	.181	.686	.314	1.000
Active Total	1.000			.663	1.000	1.000			

a. Symmetrical normalization

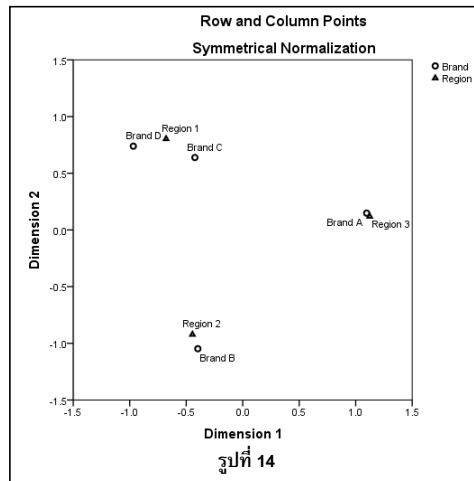
Overview Column Points ^a									
Region	Mass	Score in Dimension		Inertia	Contribution				
		1	2		Of Point to Inertia of Dimension		Of Dimension to Inertia of Point		Total
					1	2	1	2	
Region 1	.333	-.678	.803	.206	.240	.427	.476	.524	1.000
Region 2	.333	-.445	-.922	.185	.103	.563	.229	.771	1.000
Region 3	.333	1.124	.119	.272	.657	.009	.991	.009	1.000
Active Total	1.000			.663	1.000	1.000			

a. Symmetrical normalization

ပုံ ၁၂

รูปที่ 13

4. รูปที่ 14 เป็นแผนภาพการวิเคราะห์การสมนัย บางครั้งเรียกแผนภาพการรับรู้ (Perceptual Map) โดยนำ Score in Dimension ของแถวและคอลัมน์ในรูปที่ 13 มาพล็อตในแผนภาพเดียวกัน หลักในการวิเคราะห์ คือ จุดใดอยู่ใกล้ แสดงว่ามีความสอดคล้องใกล้เคียงสัมพันธ์กัน ดังนั้น จะเห็นว่าพื้นที่ 1 ใกล้ชิดกับตราสินค้า C และ D ส่วนพื้นที่ 2 มีความสัมพันธ์กับตราสินค้า B ในขณะที่พื้นที่ 3 มีความสัมพันธ์ใกล้เคียงอย่างมากกับตราสินค้า A

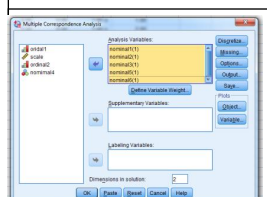
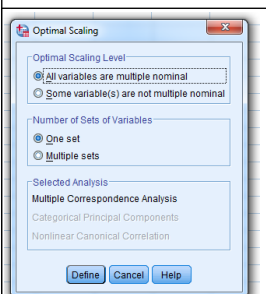
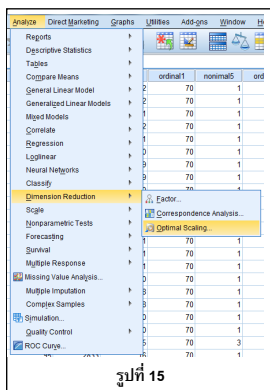


การวิเคราะห์การสมนัยเชิงพหุ

การวิเคราะห์การสมนัยเชิงพหุ (Multiple Correspondence Analysis: MCA) เป็นส่วนขยายต่อของ CA กล่าวคือ ใช้หลักการเดียวกันกับ CA เพียงแต่ข้อมูลที่ใช้ในการวิเคราะห์เป็นข้อมูลเชิงกลุ่มที่มีตั้งแต่ 3 ตัวแปรขึ้นไป โดยเป็นสถิติวิเคราะห์ตัวแปรพหุเชิงพรรณนา (Multivariate Descriptive Statistical Analysis) MCA จะทำการเปลี่ยนตัวแปรข้อมูลเชิงกลุ่ม (กรณีวัตถุประสงค์หรือประเด็นที่สนใจ) ให้เป็นตัวเลข ผลการวิเคราะห์มักนิยมแสดงในรูปของการพล็อตกราฟ โดยตัวแปรที่อยู่ในหมวดหมู่เดียวกันจะอยู่ใกล้ชิดกัน ในขณะที่ตัวแปรที่แตกต่างกันจะถูกพล็อตห่างกัน MCA ไม่มีข้อดกกลางเบื้องต้นเกี่ยวกับการแจกแจงแบบปกติของข้อมูล

ในการวิจัยเมื่อต้องแปลผลลัพธ์ของ MCA ผู้วิจัยต้องมีความเที่ยงตรงในการวิเคราะห์และไม่มีการทดสอบนัยสำคัญทางสถิติ MCA จะมีประโยชน์อย่างมากในการลดมิติของตัวแปร ทำให้เห็นความสัมพันธ์ของระดับตัวแปรได้ชัดเจนขึ้น ในบทความนี้ จะขอละรายละเอียดเกี่ยวกับ MCA ไว้ก่อนเนื่องจากต้องใช้เนื้อที่ค่อนข้างมาก ผู้ที่สนใจสามารถอ่านเพิ่มเติมได้จาก Hoffman and Leeuw (1992); Greenacre, M. (2006) อย่างไรก็ตาม หากต้องการใช้ SPSS ในการวิเคราะห์ MCA SPSS (2015) ได้ให้แนวทางคร่าวๆ ดังนี้

1. เลือกเมนู Analyze >>> Dimension Reduction >>> Optimal Scaling ดังรูปที่ 15 จะปรากฏ Dialog Box ดังรูปที่ 16
2. จากรูปที่ 16 ให้คลิกเลือก All variables multiple nominal และเลือก One set จากนั้นให้คลิก Define จะได้ Dialog Box ดังรูปที่ 17 ให้นำตัวแปรข้อมูลเชิงกลุ่มอย่าง 2 ตัวแปร ไปไว้ในช่อง Analysis Variables
3. จากรูปที่ 17 เลือกคลิก OK จะได้ผลลัพธ์ของ MCA โดยในรูปที่ 17 ผู้เขียนสมมติว่าเลือก 5 ตัวแปร คือ nominal1 nominal 2 nominal 3 nominal 5 และ nominal 6 ไว้ในช่อง Analysis Variables



สรุปและข้อเสนอแนะ

โดยทั่วไป ในการวิเคราะห์ข้อมูลเชิงกลุ่มของนักวิจัยทางสังคมศาสตร์นิยมนำเสนอด้วยความถี่และร้อยละ สำหรับสถิติในขั้นต่อมา คือ การคำนวณไคสแควร์ เพื่อวิเคราะห์ความสัมพันธ์ระหว่างตัวแปร (ทดสอบความเป็นอิสระ) อย่างไรก็ตาม ผู้วิจัยสามารถใช้การทดสอบไคสแควร์ร่วมกับวิธีการ CA เพื่อให้ได้สารสนเทศการวิจัยที่ครอบคลุมขึ้น โดยมีข้อสรุปและข้อเสนอแนะดังนี้

1. วิธีการ CA ถูกมองว่าเป็นสถิติขั้นสูงและมีการวิเคราะห์ที่ยุ่งยาก ในปัจจุบันด้วยวิทยาการทางคอมพิวเตอร์ทำให้การคำนวณเป็นเรื่องง่าย เพียงเลือกใช้โปรแกรมสำเร็จที่มีฟังก์ชันของ CA ดังเช่น โปรแกรม SPSS ที่ถูกนำเสนอไว้ในบทความนี้ นอกจากนี้ยังมี โปรแกรมอื่นอีก อาทิ STATA XLSTAT R เป็นต้น สำหรับการตีความของ CA ผู้วิจัยเพียงพิจารณาจากแผนภาพทวิที่แสดงแผนภาพการรับรู้ หากตัวแปรในแผนภาพใกล้กันแสดงว่ามีความสัมพันธ์

2. วิธีการทดสอบไคสแควร์เป็นการพิจารณาความสัมพันธ์ระหว่างตัวแปร ในขณะที่ CA นั้น จะพิจารณาความสัมพันธ์ในระดับของตัวแปรด้วย ทำให้ผู้วิจัยสามารถตีความหรือวิเคราะห์ข้อมูลเพื่อใช้ในการวางกลยุทธ์หรือให้ข้อเสนอแนะการวิจัยที่เป็นประโยชน์

3. ในการวิจัยเพื่อค้นหาลักษณะประกอบหรือปัจจัยที่ผู้วิจัยไม่ทราบมาก่อนว่าข้อมูลดังกล่าวมีกี่องค์ประกอบ เป็นแต่เพียงการวิเคราะห์จากตัวแปรจำนวนมาก (ตัวแปรเชิงพหุ) หากข้อมูลเป็นข้อมูลที่ต่อเนื่อง (Continues Data) วิธีการในการค้นหา คือ การวิเคราะห์องค์ประกอบหลัก (Principle Component Analysis: PCA) หรือการวิเคราะห์ปัจจัยเชิงสำรวจ (Exploratory Factor Analysis: EFA) ในทำนองเดียวกันหากข้อมูลในการวิเคราะห์เป็นข้อมูลเชิงกลุ่ม ผู้วิจัยสามารถใช้ CA ในการค้นหาลักษณะประกอบหรือปัจจัยได้ในลักษณะเดียวกันกับ PCA และ EFA ร่วมทั้งการมีคุณสมบัติในการลดมิติของข้อมูลด้วย

4. วิธีการ CA ไม่ใช่สถิติที่เน้นการทดสอบสมมติฐาน การทดสอบสมมติฐานจะเป็นการทดสอบด้วยไคสแควร์ในกรณีที่ผลการทดสอบ พบว่า ตัวแปรไม่เป็นอิสระต่อกัน (มีความสัมพันธ์) การนำ CA มาวิเคราะห์ต่อจะทำให้เห็นภาพความสัมพันธ์

ของระดับในตัวแปรที่เด่นชัดขึ้น และง่ายต่อการนำเสนอผลการวิจัยเมื่อพล็อตด้วย แผนภาพทวิ อย่างไรก็ตาม แม้ว่าผลการทดสอบด้วยไคสแควร์จะพบว่าตัวแปรเป็นอิสระต่อกัน ในบางครั้งเมื่อนำ CA มาวิเคราะห์ก็จะทำให้ผู้วิจัยได้พบความน่าสนใจ บางประการ

เอกสารอ้างอิง

- กมล สนิทธรรม. (2549). การวิเคราะห์การสมนัยพหุคูณและการประยุกต์. วิทยานิพนธ์ปริญญาโทบริหารธุรกิจ, มหาวิทยาลัยเชียงใหม่.
- กษิต ศาสตร์วิรุณ และกฤษ จรินทร์. (2555). ปัญหาการจัดเก็บภาษีรถยนต์มือสอง ในกรุงเทพมหานครและจังหวัดชลบุรี. การประชุมวิชาการ มหาวิทยาลัยกรุงเทพประจำปี 2555 หน้า 1-30.
- ชยุตม์ ภิรมย์สมบัติ. (2557). การวิเคราะห์ความสมนัย (Correspondence Analysis) เอกสารประกอบการบรรยายหลักสูตรสถิติวิเคราะห์ขั้นสูง สำหรับการวิจัย 2. สมาคมวิจัยสังคมศาสตร์แห่งประเทศไทย คณะครุศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย 5 สิงหาคม 2557
- ณินี พานสายตา. (2555). การรับรู้ภาพลักษณ์คุณภาพบัณฑิตหลักสูตรและการเรียนการสอนของนักศึกษาที่มีต่อสถาบันอุดมศึกษาภาครัฐ. วิทยานิพนธ์หลักสูตรวิทยาศาสตรมหาบัณฑิต (สถิติประยุกต์) คณะสถิติประยุกต์สถาบันบัณฑิตพัฒนบริหารศาสตร์.
- ราชบัณฑิตยสถาน. (2558). ศัพท์บัญญัติราชบัณฑิตยสถาน (correspondence). สืบค้นเมื่อ 3 มิถุนายน 2558 จาก <http://rirs3.royin.go.th/coinages/webcoinage.php>
- วิยะดา ต้นวัฒนากุล และคณะ (2548). ความคิดเห็นของประชาชนในจังหวัด เชียงใหม่และลำพูนที่มีต่อแนวนโยบายในการกำหนดให้วันอาทิตย์ เป็นวันทำบุญที่วัดของพุทธศาสนิกชน การประชุมวิชาการสถิติประยุกต์ฯ ภาคตะวันออกเฉียงเหนือ ครั้งที่ 1, 2-4 พ.ค. หน้า 89-99.
- Armatte, M. (2008). Histoire et prehistoire de l'analyse des donnees per J.P. Benzecri: un cas de genealogie retrospective. *Electronic Journal for History of Probability and Statistics*. 4(2), 24.

- Askell-Williams, H. and Lawson, M.J. (2004). A Correspondence Analysis of Child-Care Students' and Medical Students' Knowledge about Teaching and Learning. **International Education Journal**. 5(2), 176-205.
- Beh, E.J. (2004). Simple Correspondence Analysis: A Simple Correspondence Analysis: A Bibliographic Review. **International Statistic Review**. 72(2), 257-284.
- Bendixen, M. (2003). A Practical Guide to the Use of Correspondence Analysis in Marketing Research. **Marketing Bulletin**. 14, 1-15.
- Bertin, M. and Atanassova, I. (2015). **Factorial Correspondence Analysis Applied to Citation Contexts 2nd International Workshop on Bibliometric-enhanced Information Retrieval (BIR2015)** at the 37th European Conference on Information Retrieval (ECIR-2015), At Vienna, Austria.
- Clausen, S.E. (1988). **Applied Correspondence Analysis: An Introduction**. Thousand Oaks, CA: Sage.
- Doey, L. and Kurta, J. (2011). Correspondence analysis applied to psychological research. **Tutorials in Qualitative Methods for Psychology**. 7(1), 5-14.
- Fiedler, J.A. (1996). **A Comparison of Correspondence Analysis and Discriminant Analysis-Based Maps**. http://www.populus.com/files/Comparison%20CA_DA-Maps_f_1.pdf
- Goodman, L.A. (1996). A single general method for the analysis of cross-classified data: Reconciliation and synthesis of some methods of Pearson, Yule, and Fisher, and also some methods of correspondence analysis and association analysis. **Journal of the American Statistical Association**. 91, 408-428.
- Greenacre, M.J. (2006). **From simple to multiple correspondence analysis**. In Greenacre, M. & Blasius, J. (eds), **Multiple Correspondence Analysis and Related Methods**. Chapman & Hall/CRC Press. London. 41-76.

- Greenacre, M.J. (1984). **Theory and Application of Correspondence Analysis**. London: Academic Press.
- Guttman, L. (1953). A note on Sir Cyril Burt's "Factorial analysis of qualitative data. **The British Journal of Statistical**
- Hair, J. F. Jr., Black, W. C., Babin, B. J., and Anderson, R. E. (2010). **Multivariate Data Analysis: A global perspective (7th ed.)** Upper Saddle River, NJ: Pearson Prentice Hall.
- Hill, M.O. and Gauch Jr., H.G. (1980). **Detrended correspondence analysis: an improved ordination technique**. *Vegetatio*. 42,47–58.
- Hoffman, D. L. and Franke, G. R. (1986). Correspondence Analysis: Graphical Representation of Categorical Data in Marketing Research. **Journal of Marketing Research**. 23, 213-227.
- Hoffman, D.L. and Leeuw, D.J. (1992). Interpreting Multiple Correspondence Analysis as a Multidimensional Scaling Method. **Marketing Letters**. 3, 259-272.
- Horst, P. (1935). Measuring complex attitudes. **The Journal of Social Psychology**. 6, 369–375.
- Redfern, N. (2012). Correspondence analysis of genre preferences in UK film audiences. **Journal of Audience & Reception Studies**. 9(2), 45-55.
- Richardson M. and Kuder G. F. (1933). Making a rating scale that measures. **Personnel Journal**. 12, 36 - 40.
- SPSS. (2015). **IBM SPSS Categories 22 (Chapter 5. Correspondence Analysis)**. http://www.cc.uoa.gr/fileadmin/cc.uoa.gr/uploads/files/manuals/SPSS22/IBM_SPSS_Categories.pdf
- Stevens, S.S. (1946). **On the Theory of Scales of Measurement**. *Science*. 103, 677 - 680.
- Wei, F., Pillai, V., and Maleku, A. (2014). Sanitation in India: Role of Women's Education. **Health Science Journal**. 8(1), 90-101.

ตารางที่ 1 ตารางการนับจรรสองทางขนาด $r \times c$

แถว	คอลัมน์					รวม
	1	2	3	...	c	
1	n_{11}	n_{12}	n_{13}	...	n_{1c}	$n_{1\cdot}$
2	n_{21}	n_{22}	n_{23}	...	n_{2c}	$n_{2\cdot}$
3	n_{31}	n_{32}	n_{33}	...	n_{3c}	$n_{3\cdot}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
r	n_{r1}	n_{r2}	n_{r3}	...	n_{rc}	$n_{r\cdot}$
รวม	$n_{\cdot 1}$	$n_{\cdot 2}$	$n_{\cdot 3}$	...	$n_{\cdot c}$	n

ตารางที่ 2 Row Profile

แถว	คอลัมน์					Active Margin
	1	2	3	...	c	
1	$n_{11}/n_{1\cdot}$	$n_{12}/n_{1\cdot}$	$n_{13}/n_{1\cdot}$	...	$n_{1c}/n_{1\cdot}$	1
2	$n_{21}/n_{2\cdot}$	$n_{22}/n_{2\cdot}$	$n_{23}/n_{2\cdot}$	...	$n_{2c}/n_{2\cdot}$	1
3	$n_{31}/n_{3\cdot}$	$n_{32}/n_{3\cdot}$	$n_{33}/n_{3\cdot}$	...	$n_{3c}/n_{3\cdot}$	1
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
r	$n_{r1}/n_{r\cdot}$	$n_{r2}/n_{r\cdot}$	$n_{r3}/n_{r\cdot}$	...	$n_{rc}/n_{r\cdot}$	1
Column Mass	$n_{\cdot 1}/n$	$n_{\cdot 2}/n$	$n_{\cdot 3}/n$	...	$n_{\cdot c}/n$	1

ตารางที่ 3
 Column Profile

แถว	คอลัมน์					Row Mass
	1	2	3	...	c	
1	$n_{11}/n_{.1}$	$n_{12}/n_{.2}$	$n_{13}/n_{.3}$	...	$n_{1c}/n_{.c}$	$n_{1.}/n$
2	$n_{21}/n_{.1}$	$n_{22}/n_{.2}$	$n_{23}/n_{.3}$	...	$n_{2c}/n_{.c}$	$n_{2.}/n$
3	$n_{31}/n_{.1}$	$n_{32}/n_{.2}$	$n_{33}/n_{.3}$	...	$n_{3c}/n_{.c}$	$n_{3.}/n$
⋮	⋮	⋮	⋮	⋮	⋮	⋮
r	$n_{r1}/n_{.1}$	$n_{r2}/n_{.2}$	$n_{r3}/n_{.3}$	...	$n_{rc}/n_{.c}$	$n_{r.}/n$
Active Margin	1	1	1	...	1	1

ตารางที่ 4
 ตารางการกระจายของตราสินค้าและพื้นที่ของผู้ใช้ยาเสพติด

ตราสินค้า	พื้นที่ 1	พื้นที่ 2	พื้นที่ 3	รวม
A	5	5	30	40
B	5	25	5	35
C	15	5	5	25
D	15	5	0	20
รวม	40	40	40	120