

อันเนื่องมาแต่สูตรของยามาเน*

ABSTRACT

Thai students usually use 'Yamane's formula' for calculating the sample size, even when the students' measurement scales are not dichotomous variables. Yamane simplified the sample size formula using dichotomous variables that can be transformed proportionally, not continuously. In the proportional case the maximum variance of the sample will be 0.25, which is different from the case of continuous variables, when the maximum variance is not limited to any number. Yamane chose to simplify the original proportional sample size formula in a finite population by using 0.25, the maximum variance of the proportional case, thereby yielding 'Yamane's formula'. This new formula is a function of the population size with a random sampling error where the variance has disappeared from the formula. If the students have only two choices in their measurement scales, using 'Yamane's formula' is acceptable. However, more often the graduate students will use the five point likert scales in their research papers where scales are treated as continuous variables or metric variables. Therefore the sample size formula from the continuous variable case should be used instead of 'Yamane's formula'

*ผู้เขียนขอขอบคุณ Mr.Barry Satz ที่กรุณาตรวจแก้ภาษาอังกฤษ ขอขอบคุณ ดร.นิติ รัตนปรีชาเวช และผู้ประเมินบทความของวารสารบริหารธุรกิจ
กรุณาช่วยงานและให้ข้อเสนอแนะในการปรับปรุงข้อเขียนนี้ให้มีความสมบูรณ์ยิ่งขึ้น ข้อผิดพลาดทั้งหมดที่เกิดขึ้นในบทความ ผู้เขียนขอน้อมรับไว้แต่
เพียงผู้เดียว

บทคัดย่อ

นักศึกษาไทยนิยมใช้ “สูตรของยามานะ” ในการคำนวณหาขนาดของตัวอย่างทั้งที่มาตรวัดที่นักศึกษาใช้ไม่ใช่มาตรวัดแบบที่มีแค่สองตัวเลือก ในการสร้างสูตรของยามานะที่นักศึกษาไทยชอบใช้ยามานะเลือกแปลงมาจากสูตรในการหาขนาดของตัวอย่างที่มาตรวัดมีสองตัวเลือกหรือแบบสัดส่วน ไม่ใช่สูตรที่ใช้กับตัวแปรที่มีมาตรวัดแบบต่อเนื่อง ทั้งนี้ ค่าความแปรปรวนที่มากที่สุดของตัวอย่างกรณีของสัดส่วนสามารถแทนค่าด้วยค่า 0.25 ซึ่งต่างจากกรณีที่มีมาตรวัดแบบต่อเนื่องที่ไม่อาจระบุได้แน่ชัดว่าค่าความแปรปรวนที่มากที่สุดควรเป็นเท่าใด ยามานะเลือกแปลงสูตรในการหาขนาดของตัวอย่างกรณีสัดส่วนโดยสมมติให้ประชากรมีขนาดเล็ก และเลือกที่จะแทนค่าความแปรปรวนด้วย 0.25 ซึ่งถือว่าให้ค่าความแปรปรวนที่สูงสุดของกรณีสัดส่วนแล้ว ดังนั้นภายใต้สูตรของยามานะขนาดของตัวอย่างจึงขึ้นอยู่กับขนาดของประชากรและค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง โดยดูเหมือนกับว่าค่าความแปรปรวนของตัวอย่างหลุดหายไปจากสูตรของยามานะ อันที่จริงถ้ามาตรวัดที่นักศึกษาใช้เป็นแบบที่มีแค่สองตัวเลือกการจะใช้สูตรของยามานะก็ได้ผลดีอยู่แล้ว แต่นักวิจัยของนักศึกษานิพนธ์ส่วนใหญ่ซึ่งวัดตัวแปรที่ไม่สามารถสังเกตได้นิยมใช้มาตรวัดลิเคิร์ทแบบห้าตัวเลือกซึ่งในทางสถิติแล้วมาตรวัดนี้เป็นตัวแปรแบบต่อเนื่องหรือแบบเมตริก สูตรของยามานะที่แปลงมาจากกรณีสัดส่วนจึงไม่เหมาะที่จะนำมาใช้ ดังนั้นหากต้องการหาขนาดของตัวอย่างแบบใช้สูตรเมื่อมีการเลือกตัวอย่างแบบสุ่มจึงต้องกลับไปใช้สูตรที่เหมาะสมกับตัวแปรแบบต่อเนื่องจึงจะถูกต้อง

บทนำ

ที่มาของข้อเขียนนี้มาจากความรำคาญระคนความเหนื่อยที่จะต้องอธิบายให้นักศึกษารุ่นแล้วรุ่นเล่าที่นิยมขอสูตรของยามาเน่ (Yamane's Formula) มาใช้โดยที่ไม่ได้คิดสักนิดว่าสูตรนี้มีที่มาที่ไปอย่างไร แค่เห็นมีคนอ้างมาใช้ก็ลอกๆ ตามเขาเท่านั้น ทั้งๆ ที่หลายครั้งมาตรวัดที่ใช้ไม่ได้เอื้อที่จะให้ใช้สูตรนี้เลย หากลองไปเปิดวิทยานิพนธ์ของนักศึกษาไทยดูจะพบการอ้างสูตรนี้มาใช้บ่อยมาก ครั้งแรกๆ ผู้เขียนก็พยายามอธิบายให้นักศึกษาที่บังเอิญมาเป็นนักศึกษาในที่ปรึกษาของผู้เขียนเข้าใจแม้ว่าจะต้องใช้เวลาตามสมควร แต่รู้สึกว่าการศึกษารุ่นต่อรุ่นก็มีความเข้าใจผิดเหมือนกัน กันซะทุกรุ่น จนผู้เขียนเบื่อกับที่จะต้องมานั่งบรรยายให้นักศึกษารุ่นแล้วรุ่นเล่าฟังซ้ำๆ กันเป็นรายคน นักศึกษาคงไม่รู้สึกละอะไร แต่สำหรับคนสอนที่ต้องพูดอะไรซ้ำซากให้นักศึกษาฟังทีละคน มันสุดจะกลักร้าใจเต็มทน จึงต้องพยายามนำมาเรียบเรียงเขียนให้เห็น เพื่อให้นักศึกษาหยิบสูตรนี้มาใช้ด้วยความระมัดระวัง เพื่อโอกาสหน้าถ้ามีนักศึกษาในที่ปรึกษาหยิบสูตรนี้มาใช้ผิดๆ อีกจะได้โยนข้อเขียนนี้ให้ไปอ่านก่อนเพื่อประหยัดเวลาในการอธิบายให้สั้นลง

แต่เดิมนั้น ผู้เขียนคิดว่าคุณยามาเน่เธอไม่ได้รู้เห็นด้วยกับสูตรของเธอที่มีคนนำไปอ้างผิดอยู่บ่อยๆ โดยไม่ได้สนใจกลับไปเปิดตำราสถิติที่เธออุตส่าห์เขียนออกมาเล่มใหญ่ ว่าสูตรของเธอมันมีที่มาที่ไปอย่างไร และมีข้อสมมติอะไรบ้าง ก่อนที่จะวกเข้าสูตรของยามาเน่คงต้องทำความเข้าใจสูตรที่ใช้ในการหาขนาดของตัวอย่างทั่วๆ ไปก่อน ซึ่งเป็นสิ่งที่งัดได้จากตำราวิจัยตลาดภาษาอังกฤษแทบทุกเล่มซึ่งมักจะมีการพิสูจน์สูตรในการหาขนาดของตัวอย่างไว้ใหญ่ แต่ไม่เข้าใจว่าทำไมนักศึกษาไทยไม่นิยมใช้ ถ้าจะให้เดาเดียงเป็นเพราะไม่ทราบว่าจะหาค่าความแปรปรวน (Variance) ของตัวแปรได้อย่างไร เพราะยังไม่ได้มีการเก็บข้อมูลจริง หากเป็นตัวแปรที่ไม่รู้ว่ามีมาตรวัดเป็นเช่นไรก็คงจะจริง แต่หาประสบการณ์ที่ได้ให้คำปรึกษานักศึกษาปริญญาโทหลายๆ รุ่นพบว่า งานของนักศึกษาปริญญาโทมักจะวัดด้วยมาตรวัดลิเคิร์ต (Likert Scale) ซึ่งเป็นมาตรวัดแบบ 5 จุด เริ่มจากไม่เห็นด้วยอย่างยิ่ง (1) ถึงเห็นด้วยอย่างยิ่ง (5) เป็นส่วนใหญ่ จึงไม่น่ายากที่จะลองเดาค่าตัวเฉลี่ย ($\bar{X}$) และค่าความแปรปรวน (s^2) ของกลุ่มตัวอย่างกรณีค่าอยู่ระหว่างเท่าไรจึงจะสมเหตุผล การที่นักศึกษาพยายามหลีกเลี่ยงที่จะหาค่าความแปรปรวนเพราะไม่อยากหาค่าความแปรปรวนปรากฏให้เห็นอยู่ในสูตรเป็นทาง

เลือกแทน ความจริงแล้วการหาค่าความแปรปรวนของตัวแปรที่ใช้มาตรวัดแบบลิเคิร์ตเป็นเรื่องง่ายกว่าที่คาดคิด ไม่ได้ยากสักซักระยะจนต้องพยายามเลี่ยงไปใช้สูตรของยามาเน่ อย่างไรก็ตามก่อนที่จะสงสัยว่าทำไมผู้เขียนจึงดูหมิ่นสูตรของยามาเน่ ลองมาดูที่มาของสูตรที่ใช้ในการหาขนาดของตัวอย่างที่นิยมใช้กันโดยทั่วไปก่อนดีกว่า

จากการที่ใช้ตำราการวิจัยตลาดภาษาอังกฤษของผู้แต่งหลายๆ ท่าน สิ่งที่ผู้เขียนพอจะได้จากตำราที่ใช้ทุกเล่มว่าจะมีวิธีการพิสูจน์สูตรที่ใช้ในการหาขนาดของตัวอย่างคล้ายคลึงกัน และผู้เขียนก็เพิ่งพบว่าท่านใดพูดถึงสูตรของยามาเน่เลย แต่กลับพบว่าตำราวิจัยตลาดของผู้แต่งชาวไทยบางเล่มมักจะอ้างถึงสูตรนี้ โดยผู้เขียนเองก็ไม่แน่ใจว่าคุณยามาเน่เธอรู้เรื่อง การอ้างชื่อเธอผิดๆ หรือไม่ ซึ่งการที่จะได้ขนาดของตัวอย่างที่เหมาะสมจะทำให้หลายวิธี อย่างไรก็ตามการใช้สูตรจะทำได้ก็ต่อเมื่อเงื่อนไขของตัวแปรมาจากการสุ่ม (Random) หรือสุ่มแบบสุ่มสี่สุ่มห้าสูตรที่ใช้ก็คงให้ค่าที่ไม่ถูกต้องเพราะหากไม่ใช้การเลือกตัวอย่างแบบสุ่มหรือที่นักศึกษาที่เรียนวิจัยตลาดรู้จักกันว่า การเลือกตัวอย่างแบบใช้ความน่าจะเป็น (Probability Sampling) ก็จะไม่สามารถหาขนาดของตัวอย่างที่ให้ค่าระดับความเชื่อมั่นที่ต้องการโดยยอมให้ความคลาดเคลื่อนจากการสุ่มตัวอย่างอยู่ในระดับที่ยอมรับได้ด้วย

ก่อนอื่นต้องทำความเข้าใจการแจกแจงของตัวแปรทางสถิติ ซึ่งมีอยู่มากมายหลายแบบ แต่เท่าที่ผู้เขียนเคยเห็นการอธิบายถึงที่มาของขนาดของตัวอย่างมักจะใช้ 2 แบบ คือตัวแปรที่มีการแจกแจงแบบปกติ (Normal Distribution) และตัวแปรที่มีการแจกแจงแบบ 2 ค่า (Binomial Distribution) ในทางวิจัยตลาดจะเรียกตัวแปรประเภทหลังนี้ว่าตัวแปรที่มี 2 ค่า (Dichotomous Variable) สำหรับตัวแปรที่จะใช้การแจกแจงแบบปกติได้จะต้องเป็นตัวแปรที่มีมาตรวัดแบบต่อเนื่อง (Continuous Variable) ซึ่งก็กินความถึงตัวแปรที่มีมาตรวัดแบบช่วง (Interval Scale) และตัวแปรที่มีมาตรวัดแบบสัดส่วน (Ratio Scale) เท่านั้น สำหรับตัวแปรที่มีมาตรวัดแบบแบ่งกลุ่ม (Nominal Scale) และที่มีมาตรวัดแบบเรียงลำดับ (Ordinal Scale) ไม่อาจจัดอยู่ในตัวแปรที่มีมาตรวัดแบบต่อเนื่อง (Continuous Variable) ได้ ยกเว้นตัวแปรนั้นมีแค่ 2 ค่า หรือที่สามารถแปลงให้อยู่ในรูปสัดส่วน (Proportion) ได้

ในกรณีของสัดส่วน ถ้าให้

โอกาสที่จะประสบความสำเร็จ

และ โอกาสที่จะไม่ประสบความสำเร็จ

$$= \pi$$

$$= 1-\pi$$

นั่นหมายความว่าจากการสุ่มตัวอย่างแต่ละครั้งสามารถหาสัดส่วนของโอกาสที่จะประสบความสำเร็จจากโอกาสทั้งหมดซึ่งเท่ากับหนึ่งได้ ในกรณีเช่นนี้มาตรวัดจะมีแค่ 2 ค่า คือ ถ้าใช่ มาตรวัด = 1 ถ้าไม่ใช่ มาตรวัด = 0 หรือที่รู้จักกันในทางสถิติคือการใช้ตัวแปรหุ่น (Dummy Variable) ด้วยการลงรหัสค่าทั้งสองของตัวแปรให้เท่ากับหนึ่งและศูนย์เท่านั้น ในกรณีเช่นนี้แม้ว่ามาตรวัดจะเป็นมาตรวัดแบบไม่ต่อเนื่องก็ตาม แต่การแทนค่ามาตรวัดด้วยหนึ่งและศูนย์จะสามารถหาค่าถ้อยเฉลี่ยและค่าความแปรปรวนของตัวแปรได้เหมือนกับตัวแปรที่มีมาตรวัดแบบต่อเนื่องเช่นกัน ทั้งนี้ค่าถ้อยเฉลี่ยของตัวแปรจะเท่ากับสัดส่วนของจำนวนผู้เลือกทางเลือกที่ใช้และลงรหัสเป็นหนึ่งต่อจำนวนตัวอย่างทั้งสิ้น จะทำให้ได้สัดส่วนซึ่งเป็นตัวเลขที่เปรียบเสมือนต่อเนื่อง จึงหาค่าประมาณการช่วง (Interval Estimate) ของสัดส่วนของประชากรได้เช่นเดียวกับตัวแปรแบบต่อเนื่องนั่นเอง ดังนั้นหากกำหนดระดับความเชื่อมั่นที่ต้องการและความคลาดเคลื่อนจากการสุ่มตัวอย่างที่พอรับได้ ก็จะนำไปสู่การกำหนดขนาดของตัวอย่างที่เหมาะสมได้

การกระจายของการสุ่มตัวอย่าง กรณีตัวแปรต่อเนื่อง

ที่มาของทฤษฎีทางสถิติที่นักวิจัยใช้เป็นพื้นฐานในการศึกษาสถิติเบื้องต้น ได้แก่ Central Limit Theorem



รูปที่ 1: การกระจายของประชากรรูปแบบต่างๆ ทั้งแบบปกติและไม่ปกติ

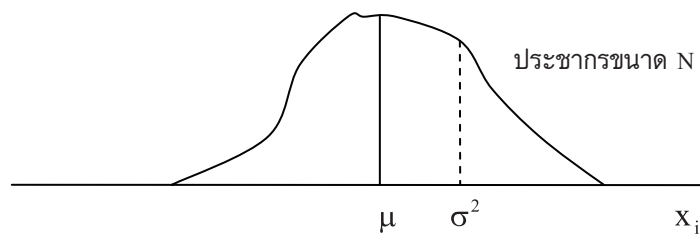
ตามทฤษฎีนี้ได้มาจากการวิจัยของนักสถิติในสมมติฐานที่ถือว่ากลุ่มตัวอย่าง¹ ที่มีขนาดตัวอย่างเท่ากันหลายๆ ครั้ง แต่ละครั้งก็จะได้ค่าถ้อยเฉลี่ยของตัวอย่าง ($\bar{X}$) มาหนึ่งตัว และพบว่าเมื่อขนาดของตัวอย่าง (n) ใหญ่ขึ้นจะทำให้การกระจายของค่าถ้อยเฉลี่ยของตัวอย่าง ($\bar{X}$) จากการสุ่มตัวอย่างเข้าใกล้การแจกแจงแบบปกติ (Normal Distribution) ไม่ว่าประชากรจะมีการแจกแจงแบบใดก็ตาม โดยค่าที่คาดไว้ (Expected Value) ของค่าถ้อยเฉลี่ยของตัวอย่าง ($E(\bar{X})$) เท่ากับค่าถ้อยเฉลี่ยของประชากร (μ) และมีค่าความแปรปรวนของค่าถ้อยเฉลี่ยของตัวอย่าง ($\sigma_{\bar{X}}^2$) เท่ากับค่าความแปรปรวนของประชากร (σ^2) หารด้วยขนาดของตัวอย่าง (n)

เพื่อความเข้าใจที่ตรงกันขอแยกประเภทของการกระจายไว้ก่อน ดังนี้

1. การกระจายของประชากร (Population) ขนาด N ตาม Central Limit Theorem จะมีการกระจายแบบใดก็ได้ขึ้นอยู่กับต่อไปนี้ (Zinkmund, 2003, p.459)

ในทางสถิติอาจสรุปภาพรวมของประชากรไม่ว่าประชากรจะมีการกระจายแบบใดด้วยค่าบางค่าที่เรียกว่า Population Parameters ซึ่งสามารถอธิบายลักษณะภาพรวมของประชากรได้อย่างกะทัดรัด โดยทั่วไปจะมีอาทิเช่น ค่าถ้อยเฉลี่ยของประชากร (μ) และค่าความแปรปรวนของประชากร (σ^2) เป็นต้น

¹ สุ่มตัวอย่างนี้หมายความว่า สุ่มที่มาจากภาษาอังกฤษคือ Random ไม่ใช่การพูดถึงการเลือกตัวอย่างโดยทั่วไปซึ่งมีทั้งที่สุ่มและไม่สุ่มแต่ก็มีผู้ใช้คำรวมๆ กันว่าสุ่มตัวอย่าง แทนที่จะเรียกว่าการเลือกตัวอย่างซึ่งเป็นคำกว้างๆ และดูจะเหมาะสมกว่าและสามารถใช้ครอบคลุมได้ทั้งการเลือกตัวอย่างแบบสุ่มและไม่สุ่ม



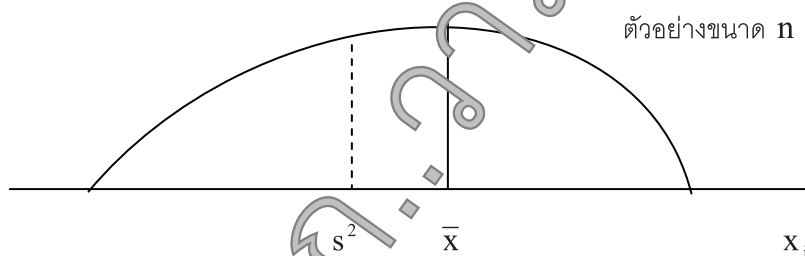
รูปที่ 2: การกระจายของประชากรและ Population Parameters

ค่าเฉลี่ยของประชากร, $\mu = \frac{\sum_{i=1}^N x_i}{N}$

และมีความแปรปรวนของประชากร, $\sigma^2 = \frac{\sum_{i=1}^N (x_i - \mu)^2}{N}$

2. การกระจายของตัวอย่าง (Sample) ขนาด n ซึ่งนักสถิตินิยมสรุปให้เห็นภาพรวมด้วยค่า Sample statistics ได้แก่ค่าเฉลี่ยของตัวอย่าง ($\bar{x}$) และค่าความ

แปรปรวนของตัวอย่าง (s^2) ทั้งนี้ตาม Central Limit Theorem ไม่บอกให้ทราบว่า การกระจายของตัวอย่างเป็นรูปร่างแบบใด



รูปที่ 3: การกระจายของกลุ่มตัวอย่างและ Sample Statistics

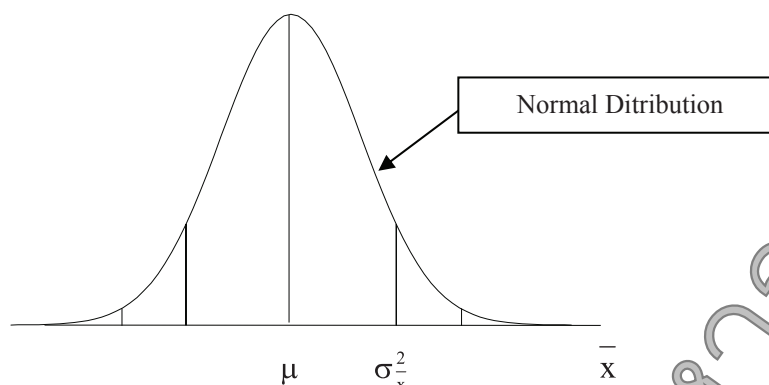
ค่าเฉลี่ยของตัวอย่าง, $\bar{x} = \frac{\sum_{i=1}^n x_i}{n}$

และค่าความแปรปรวนของตัวอย่าง, $s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1}$

3. การกระจายของค่าเฉลี่ยของตัวอย่าง (Sample Mean, $\bar{x}$) ไม่มีการเลือกตัวอย่างที่มีขนาดเท่ากันซ้ำแล้วซ้ำอีกตามที่นักสถิติได้พิสูจน์ไว้แล้ว เป็นสิ่งที่ Central Limit Theorem บอกไว้ว่า ถ้าขนาดของตัวอย่าง (n) ใหญ่พอ จะทำให้การกระจายของค่าเฉลี่ยของตัวอย่างเข้าใกล้การ

กระจายแบบปกติ (Normal Distribution) โดยมีค่าที่คาดหวัง (Expected Value)² ของค่าเฉลี่ยของตัวอย่าง ($E(\bar{x})$) เท่ากับค่าเฉลี่ยของประชากร (μ) และมีความแปรปรวนของค่าเฉลี่ยของตัวอย่างเท่ากับค่าความแปรปรวนของประชากร (σ^2) หารด้วยขนาดของตัวอย่าง (n)

² หรือ *E* เหมือนหนึ่งว่าเป็นค่าเฉลี่ยของค่าเฉลี่ยของตัวอย่างอีกทีหนึ่ง



รูปที่ 4: การกระจายของค่าเฉลี่ยของกลุ่มตัวอย่าง

ค่าที่คาดไว้ของค่าเฉลี่ยของตัวอย่าง, $E(\bar{X}) = \mu$

และมีค่าความแปรปรวนของค่าเฉลี่ยของตัวอย่าง, $\sigma^2_{\bar{x}} = \frac{\sigma^2}{n}$

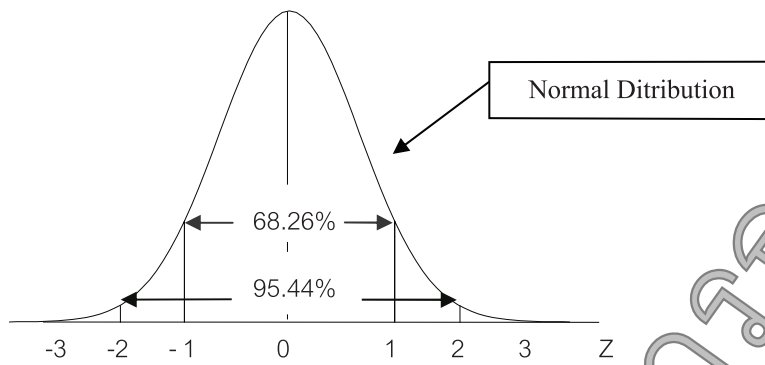
หรือ $= \frac{s^2}{n}$ ถ้าไม่ทราบค่า σ^2

จาก Central Limit Theorem นี้เอง มีข้อสังเกตว่าการที่จะมีการกระจายแบบปกติได้ ตัวแปรจะต้องเป็นตัวแปรต่อเนื่อง (Continuous Variable) และการที่ตัวแปรจะเป็นตัวแปรต่อเนื่องได้ก็หมายความว่า ณ จุดที่ $x_i = a$ จะหาค่าความน่าจะเป็นของตัวแปรไม่ได้ นั่นคือ Probability ($x_i = a$) = 0 การจะหาค่าความน่าจะเป็นของตัวแปรต่อเนื่องจะต้องหาค่าความน่าจะเป็นระหว่างค่า x_i ที่มีค่าระหว่าง 2 จุด เช่น Probability ($a < x_i < b$) = $\int_a^b f(x) dx$ ซึ่งมากกว่าศูนย์ (Lindgren, 1968, p.67) ดังนั้น หากมีตัวแปรที่ไม่ต่อเนื่อง (Discrete Variable) การที่จะบอกให้มีการกระจายแบบปกติแล้วหาค่าเฉลี่ยและค่าความแปรปรวนตามสูตรที่กล่าวมาแล้วจึงทำไม่ได้ ดังนั้น ในการหาขนาดของตัวอย่างจึงมีที่มาจาก Central Limit Theorem นี้เอง

จากความรู้เรื่องการกระจายแบบปกติ (Normal Distribution) ทำให้ทราบ ถ้าตัวแปรใดมีการกระจายแบบ

ปกติก็จะสามารถแปลงให้อยู่ในรูปของคะแนนมาตรฐาน (Z-Scores) ซึ่งมีการกระจายแบบปกติ (Standardized Normal Distribution) ซึ่งมีค่าเฉลี่ย (Mean) เท่ากับศูนย์ และค่าความคลาดเคลื่อนมาตรฐาน (Standard Deviation) เท่ากับ ± 1 เช่นเดียวกับค่าความแปรปรวน (Variance) ซึ่งเท่ากับหนึ่ง เช่นกัน และภายในช่วงของ Z-Scores ที่เท่ากับ 10 ± 1 จะมีความน่าจะเป็นคงที่ซึ่งเท่ากับ 68.26% ส่วนภายในช่วงของ Z-Scores ที่เท่ากับ ± 2 จะมีความน่าจะเป็นเท่ากับ 95.44% และภายในช่วงของ Z-Scores ที่เท่ากับ ± 3 จะมีความน่าจะเป็นที่เท่ากับ 99.73% ไม่ว่าขนาดของตัวอย่างจะเป็นเท่าไรก็ตาม ทราบใดที่ขนาดของตัวอย่างนั้นใหญ่พอและที่ใหญ่พอก็แค่ 30 ตัวอย่างขึ้นไปก็สามารถใช้ Central Limit Theorem ได้แล้ว แต่เรามักจะเห็นการใช้ระดับความน่าจะเป็นที่เป็นตัวเลขกลมๆ คือ 95% ซึ่งเป็นขนาดของพื้นที่ของความน่าจะเป็นภายในช่วงของ Z-Scores ที่เท่ากับ ± 1.96 นั่นเอง

³ ต่อจากนี้ เพื่อความสะดวกจะขอเรียกคะแนนมาตรฐานว่า Z-Scores หรือเรียกสั้นๆ ว่า Z แทนที่จะเรียกว่าคะแนนมาตรฐาน



รูปที่ 5: การกระจายแบบปกติที่แปลงให้อยู่ในรูปคะแนนมาตรฐาน

จากความรู้เรื่อง Central Limit Theorem และ Normal Distribution ทำให้ทราบว่าจากการสุ่มตัวอย่างเพียงครั้งเดียว หากขนาดของตัวอย่าง (n) มีค่าตั้งแต่ 30 ตัวอย่างขึ้นไป ค่าเฉลี่ยของตัวอย่าง ($\bar{x}$) ไม่ควรอยู่ห่างจากค่าเฉลี่ยของประชากร (μ) เกินกว่าค่า Z ที่ ± 1.96 คูณด้วยค่าความคลาดเคลื่อนมาตรฐานของค่าเฉลี่ย ($\sigma_{\bar{x}}$) ด้วยระดับความเชื่อมั่น 95%

ดังนั้น หากระบุให้ Z_{α} สัมพันธ์กับระดับความเชื่อมั่น $(1-\alpha)$ จะได้

$$\mu - Z_{\alpha} \sigma_{\bar{x}} \leq \bar{x} \leq \mu + Z_{\alpha} \sigma_{\bar{x}}$$

$$\mu - Z_{\alpha} \frac{\sigma}{\sqrt{n}} \leq \bar{x} \leq \mu + Z_{\alpha} \frac{\sigma}{\sqrt{n}}$$

$$- Z_{\alpha} \frac{\sigma}{\sqrt{n}} \leq \bar{x} - \mu \leq + Z_{\alpha} \frac{\sigma}{\sqrt{n}}$$

หรือ
$$(\bar{x} - \mu)^2 = \frac{Z_{\alpha}^2 \sigma^2}{n}$$

และ
$$n = \frac{Z_{\alpha}^2 \sigma^2}{(\bar{x} - \mu)^2}$$

เมื่อ Z_{α} เป็นค่า Z-score ซึ่งได้จากการกำหนดระดับความเชื่อมั่นที่ $(1-\alpha)$ อาทิเช่น $Z_{0.05} = 1.96$ หมายความว่ามีความเชื่อมั่น 95% ว่าค่า ($\bar{x}$) ควรจะอยู่ห่างจาก (μ) ไม่เกิน ± 1.96 เท่าของ $\sigma_{\bar{x}}$ ทั้งนี้ μ จะไม่ทราบค่าเฉลี่ยของประชากร และการสุ่มตัวอย่างจะทำเพียงครั้งเดียว ซึ่งทำให้ได้ค่าเฉลี่ยของกลุ่มตัวอย่าง ($\bar{x}$) เพียงค่าเดียว แต่จาก Central Limit Theorem ทำให้ทราบว่า 95 ครั้งจาก 100 ครั้งที่ทำการสุ่มตัวอย่าง ค่าเฉลี่ยของ

ประชากรควรตกอยู่ในช่วง $\bar{x} - (1.96) \sigma_{\bar{x}}$ ถึง $\bar{x} + (1.96) \sigma_{\bar{x}}$ นั่นเอง ซึ่งแน่นอนย่อมมีโอกาสพลาดไปประมาณ 5 ครั้งจาก 100 ครั้ง

σ^2 เป็นค่าความแปรปรวนของประชากร และหากไม่ทราบค่านี้สามารถใช้ค่าความแปรปรวนของตัวอย่าง (s^2) แทนที่ได้

n เป็นขนาดของตัวอย่างซึ่งอย่างน้อยต้องไม่ต่ำกว่า 30 ตัวอย่าง เพื่อให้ค่าเฉลี่ยของตัวอย่างมีการกระจายแบบปกติ

$(\bar{x} - \mu)$ เป็นค่าผลต่างระหว่างค่าเฉลี่ยของตัวอย่างซึ่งเกิดขึ้นจากการสุ่มตัวอย่าง กับค่าเฉลี่ยของประชากร เพราะหากเลือกตัวอย่างจากทั้งหมดของประชากร โดยที่การวัดทำอย่างถูกต้องค่า $(\bar{x} - \mu)$ จะต้องเท่ากับศูนย์ ทั้งนี้เพราะจะได้ค่าเฉลี่ยของประชากร (μ) แทนที่จะเป็นค่าเฉลี่ยของตัวอย่าง ($\bar{x}$) แต่โดยทั่วไปขนาดของตัวอย่างจะน้อยกว่าขนาดของ ประชากรทำให้ค่าทั้งสองต่างกัน จึงอาจเรียกค่านี้ว่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (Random Sampling Error) หรือในตำราบางเล่ม (Malhotra, 2009, p.403) เรียกค่านี้ว่าระดับความแม่นยำ (Precision Level) นั่นเอง ถ้าค่าความแม่นยำสูง แสดงว่าค่าความคลาดเคลื่อนจากการสุ่มตัวอย่างมีขนาดเล็ก

และถ้าแทนค่า $(\bar{x} - \mu) = e$ ซึ่งได้แก่ความคลาดเคลื่อนจากการสุ่มตัวอย่าง จะได้สูตรในการหาขนาดของตัวอย่างในกรณีประชากรมีขนาดใหญ่ (Infinite Population) ดังนี้

$$n = \frac{Z_{\alpha}^2 \sigma^2}{e^2}$$

(1)

จากสูตรนี้ขนาดของตัวอย่างจะขึ้นอยู่กับค่า Z_{α} ที่สัมพันธ์กับระดับความเชื่อมั่น $(1-\alpha)$ และค่าความแปรปรวนของประชากร (σ^2) ในทางเดียวกัน แต่ขนาดของตัวอย่างจะสัมพันธ์ในทางผกผันกับค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) ซึ่งคงไม่ใช่ว่ามีแต่ยามานเท่านั้นที่จะได้สูตรนี้ ตำราวิจัยตลาดเท่าที่ผู้เขียนเคยใช้มาทุกเล่มก็ได้สูตรนี้เช่นกัน เป็นสูตรตั้งต้น และเชื่อว่านักศึกษาที่เรียนสถิติเบื้องต้นมาแล้วคงพิสูจนสูตรนี้ได้ไม่ยาก

ภายใต้สูตรนี้มีสมมติฐานว่าประชากรมีขนาดใหญ่ (Infinite Population) ซึ่งในทางปฏิบัติแค่ 500,000 ตัวอย่างขึ้นไปก็น่าจะถือว่าใหญ่ (Infinity) ได้แล้วแน่นอนเมื่อประชากรมีขนาดใหญ่ควรจะใช้สูตรนี้ แต่เมื่อประชากรมีขนาดเล็กจะต้องปรับสูตรอีกเล็กน้อย ในกรณีที่ประชากรมีขนาดเล็กค่าความแปรปรวนจะเล็กลงไปด้วย แต่ไม่เล็กลงตามสัดส่วนของประชากรที่ลดลง การคำนวณค่าความแปรปรวนของประชากรที่มีขนาดเล็กลง (Finite Population) จะต้องปรับด้วยตัวปรับค่าประชากรขนาดเล็ก (Finite Population Correction Factor; FPC) ซึ่งได้แก่ $\frac{N-n}{N-1}$ เมื่อ N คือขนาดของประชากร และ n คือขนาดของตัวอย่าง ซึ่งแน่นอนว่า $\frac{N-n}{N-1}$ จะต้องน้อยกว่าหนึ่งตราบดีที่ n มากกว่าหนึ่ง นั่นคือความแปรปรวนของประชากรที่มีขนาดเล็กจะเล็กกว่าความแปรปรวนของประชากรที่มีขนาดใหญ่ด้วยอัตราส่วน $\frac{N-n}{N-1}$ นั่นเอง

ดังนั้นในกรณีนี้ ค่าความแปรปรวนของประชากรที่มีขนาดเล็ก (Finite Population) จะเป็น $\sigma^2 \cdot \frac{N-n}{N-1}$

และค่าความแปรปรวนของค่าเฉลี่ยของตัวอย่าง

$$\sigma_x^2 = \sigma^2 \cdot \frac{N-n}{N-1}$$

ค่าคลาดเคลื่อนของความคลาดเคลื่อนจากการสุ่มตัวอย่าง จะเป็น

$$(\bar{x} - \mu)^2 \text{ หรือ } e^2 = \frac{Z_{\alpha}^2 \sigma^2}{n} \cdot \frac{N-n}{N-1}$$

$$n(N-1)e^2 = Z_{\alpha}^2 \sigma^2 N - Z_{\alpha}^2 \sigma^2 n$$

$$n[(N-1)e^2 + Z_{\alpha}^2 \sigma^2] = Z_{\alpha}^2 \sigma^2 N$$

ดังนั้นในกรณีประชากรที่มีขนาดเล็ก (Finite Population) ขนาดของตัวอย่างจะเป็น

$$n = \frac{Z_{\alpha}^2 \sigma^2 N}{(N-1)e^2 + Z_{\alpha}^2 \sigma^2} \quad (2)$$

ซึ่งหมายความว่าค่าคลาดเคลื่อนจากค่าระดับความเชื่อมั่น $(1-\alpha)$ ซึ่งเป็นตัวกำหนดค่าความแปรปรวนของประชากร Z_{α} และค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) ซึ่งอยู่ทางด้านขวาของสมการที่ (1) จะมีอิทธิพลในการกำหนดขนาดของตัวอย่างแล้วขนาดของประชากร (N) จะมีส่วนในการกำหนดขนาดของตัวอย่างในกรณีที่ประชากรมีขนาดเล็กด้วย

ตามสูตรข้างต้นก็ยังไม่เข้าถึงสูตรของยามานที่นักศึกษานักวิจัยไทยชอบใช้ เพราะในกรณีของสูตรข้างต้นมาตรวัดจำเป็นต้องมีแค่ 2 ค่า จะมีค่าเป็นอย่างไก็ได้ไม่จำกัดหน่วยที่ใช้ในการวัด แต่จะต้องเป็นมาตรวัดแบบช่วง (Interval Scale) หรือมาตรวัดแบบสัดส่วน (Ratio Scale) เท่านั้น เช่นเดียวกับค่าความแปรปรวนของประชากรไม่ได้จำกัดว่าจะเกิน 0.25 ไม่ได้ดังเช่นการกระจายของตัวแปรแบบสัดส่วน (Proportion)

การกระจายของการสุ่มตัวอย่าง

กรณีสัดส่วน

ในกรณีที่มาตรวัดมีแค่ 2 ค่า สามารถแสดงให้เห็นในรูปของสัดส่วนซึ่งมีค่าความน่าจะเป็นรวมกันเท่ากับหนึ่ง ดังนั้นแทนที่ค่าเฉลี่ยของตัวอย่างจะเป็น $\bar{x}$ แต่จะมีค่าเท่ากับค่าของสัดส่วนของตัวอย่างซึ่งเท่ากับ p แทนจาก Central Limit Theorem ค่าการกระจายของสัดส่วนของตัวอย่าง (p) สำหรับตัวอย่างขนาด n จะมีค่าที่คาดไว้ของสัดส่วนของตัวอย่าง $(E(p))$ และค่าความแปรปรวนของสัดส่วนของตัวอย่าง (σ_p^2) ดังนี้

อันเนื่องมาแต่สูตรของยามานะ

$$E(p) = \pi$$

$$\sigma_p^2 = \frac{\pi(1-\pi)}{n} \quad \text{กรณีประชากรมีขนาดใหญ่ (Infinite Population)}$$

หรือ $\sigma_p^2 = \frac{\pi(1-\pi)}{n} \frac{N-n}{N-1} \quad \text{กรณีประชากรมีขนาดเล็ก (Finite Population)}$

เมื่อ π คือ สัดส่วนของประชากร

σ_p^2 คือ ความแปรปรวนของสัดส่วนของตัวอย่าง

n คือ ขนาดของตัวอย่าง

N คือ ขนาดของประชากร

ซึ่งสามารถหาค่าขนาดของตัวอย่างได้เช่นเดียวกัน ตัวอย่างเท่ากับ p จาก Central Limit Theorem จะสรุปได้ว่าสัดส่วนของประชากร (π) จะอยู่ระหว่างค่า $p \pm Z_\alpha \sigma_p$ กับในกรณีตัวแปรเป็นตัวแปรต่อเนื่อง (Continuous Variable) แม้จะสุ่มตัวอย่างเพียงครั้งเดียวและได้สัดส่วนของ

$$p - Z_\alpha \sigma_p \leq \pi \leq p + Z_\alpha \sigma_p$$

เมื่อ p คือ สัดส่วนของตัวอย่าง

π คือ สัดส่วนของประชากร

σ_p คือ ค่าความคลาดเคลื่อนมาตรฐานของสัดส่วนของตัวอย่าง

Z_α คือ ค่า Z-scores ณ ระดับความเชื่อมั่น $(1-\alpha)$

หากทราบค่าสัดส่วนของประชากร (π), $\sigma_p^2 = \frac{\pi(1-\pi)}{n}$

ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่างของสัดส่วน จะเท่ากับ $(p - \pi)$ หรือสมมติให้เท่ากับ e

$$\text{จะได้ } (p - \pi)^2 = e^2 = Z_\alpha^2 \sigma_p^2$$

$$e^2 = Z_\alpha^2 \frac{\pi(1-\pi)}{n}$$

$$n = \frac{Z_\alpha^2 \pi(1-\pi)}{e^2} \quad (3)$$

หรือหากไม่ทราบสัดส่วนของประชากร (π) ก็ใช้ สัดส่วนของตัวอย่าง (p) แทน

ดังนั้น $n = \frac{Z_\alpha^2 p(1-p)}{e^2} \quad \text{กรณี Infinite Population}$

หรือ $n = \frac{Z_\alpha^2 pq}{e^2} \quad \text{เมื่อ } q = 1-p$

ที่มาของสูตรของยามาเน่

จาก $e^2 = \frac{Z_\alpha^2 \pi(1-\pi)}{n}$ ในกรณีของสัดส่วนเมื่อประชากรมีขนาดใหญ่

ถ้าประชากรมีขนาดเล็ก (Finite Population) ค่าความแปรปรวนของสัดส่วนของตัวอย่างจะปรับตัวปรับค่าประชากรขนาดเล็ก (Finite Population Correction Factor; FPC) สมการข้างต้นควรจะเป็น

$$e^2 = \frac{Z_\alpha^2 \pi(1-\pi)}{n} \cdot \frac{N-n}{N-1} \quad (4)$$

แต่ครั้งนี้ แทนที่ยามาเน่จะใช้ FPC ซึ่งเท่ากับ $\frac{N-n}{N-1}$ ยามาเน่กลับใช้ $\frac{N-n}{N}$ แทน ทั้งๆ ที่ก่อนหน้านี้ยามาเน่ใช้ $\frac{N-n}{N-1}$ ทั้งนี้อาจเป็นเพราะในกรณีที่ขนาดของประชากร N มีขนาดใหญ่มากสัดส่วนทั้งสองแทบจะเหมือนกัน จึงมีผลให้ยามาเน่ได้สมการนี้แทน

$$\begin{aligned} e^2 &= \frac{Z_\alpha^2 \pi(1-\pi)}{n} \cdot \frac{N-n}{N} \quad (\text{Yamane, 1973, p.729}) \quad (5) \\ nNe^2 &= Z_\alpha^2 \pi(1-\pi)N - Z_\alpha^2 \pi(1-\pi)n \\ n[Ne^2 + Z_\alpha^2 \pi(1-\pi)] &= Z_\alpha^2 \pi(1-\pi)N \\ \text{และ } n &= \frac{Z_\alpha^2 \pi(1-\pi)N}{Z_\alpha^2 \pi(1-\pi) + Ne^2} \end{aligned}$$

และในกรณีของสัดส่วนค่า σ_p^2 จะสูงสุดเมื่อ $\pi = 0.5$ และ $1 - \pi = 0.5$ ซึ่งจะทำให้ค่า n สูงสุดเช่นกัน

ยามาเน่แทนค่า $\pi = 0.5$ และ $1 - \pi = 0.5$

$$\begin{aligned} \text{ค่า } n &= \frac{Z_\alpha^2 (0.5)^2 N}{Z_\alpha^2 (0.5)^2 + Ne^2} \\ \text{เมื่อ } n &\text{ คือ ขนาดของตัวอย่าง} \\ N &\text{ คือ ขนาดของประชากร} \\ e &\text{ คือ ความคลาดเคลื่อนจากการสุ่มตัวอย่าง} \\ \text{และถ้าให้ } Z &= 1.96 \text{ คือ ค่า Z-scores ณ ระดับความเชื่อมั่น 95\%} \\ \text{ดังนั้น } n &= \frac{(1.96)^2 (0.5)^2 N}{(1.96)^2 (0.5)^2 + Ne^2} \end{aligned}$$

สำหรับค่า Z ซึ่งเท่ากับ 1.96 ยามาเน่ ใช้เลขตัวกลม คือ 2 แทน จะได้

$$n = \frac{(2)^2 (0.5)^2 N}{(2)^2 (0.5)^2 + Ne^2}$$

หรือ

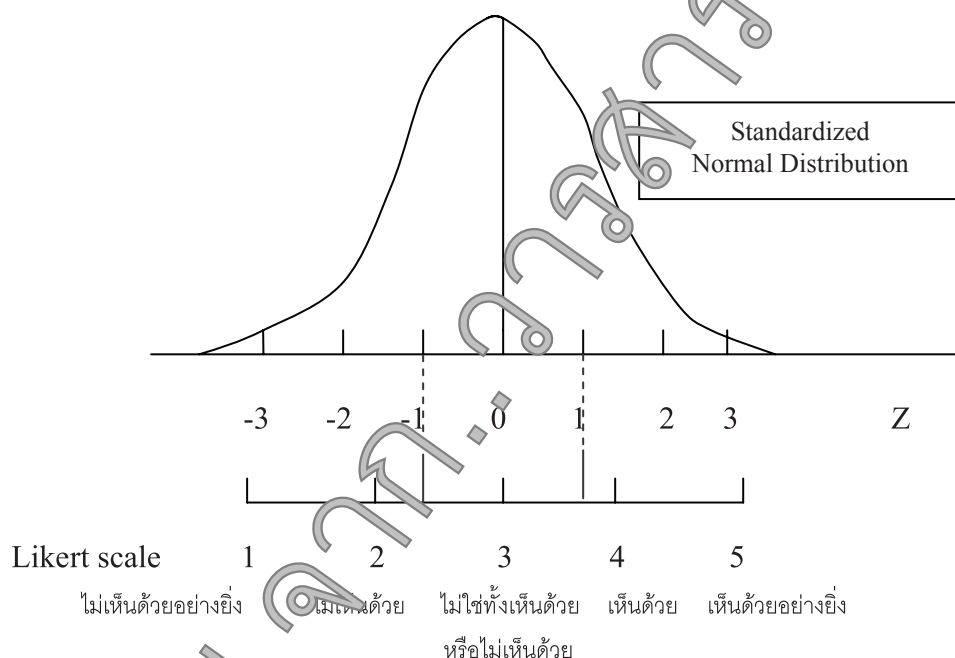
$$n = \frac{N}{1 + Ne^2}$$

(6)

อันเนื่องมาแต่สูตรของยามานะ

นั่นหมายความว่ายามานะมีข้อสมมติอยู่ 2 ประการ คือ สมมติว่าสัดส่วน (Proportion) เท่ากับ 0.5 ซึ่งเป็นค่าที่ทำให้ σ_p^2 มีค่าสูงสุดที่ 0.25 และขนาดของตัวอย่าง (n) มีค่าสูงสุด และหากสมมติระดับความเชื่อมั่นที่ 95.44% ซึ่งให้ค่า Z เท่ากับ 2 จะได้สูตรที่นักศึกษาไทยชอบใช้แล้วเรียกว่า สูตรของยามานะ ซึ่งได้แก่สมการที่ (6) ทั้งๆ ที่มาตรวัดที่ท่านเหล่านั้นใช้ไม่ได้มีแค่ 2 ค่า จึงไม่อาจแปลงเป็นค่าสัดส่วน (Proportion) ได้ โดยเฉพาะนักศึกษาปริญญาโทที่นิยมวัดตัวแปรประเภทที่ไม่สามารถสังเกตได้ (Unobserved หรือ Latent Variable) โดยใช้มาตรวัดลิเคิร์ตแบบห้าตัวเลือก (5 Point Likert Scale) เป็นส่วนใหญ่ หรือใช้มาตรวัดที่เป็นตัวแปรต่อเนื่องที่มีมากกว่า 2 ค่า หรือวัตถุประสงคฺของงานวิจัยไม่ได้ต้องการพยากรณ์สัดส่วนของตัวแปรซึ่งมีแค่ 2

ค่าแต่อย่างใด ผู้เขียนจึงไม่เข้าใจว่าท่านเหล่านั้นจะนำสูตรของยามานะนี้มาใช้ได้อย่างไร เมื่อพิจารณาว่ามาตรวัดลิเคิร์ตแบบห้าตัวเลือกค่าต่ำสุดที่เป็นไปได้จะเป็น (1) และค่าสูงสุดที่เป็นไปได้จะเป็น (5) หากลองสมมติต่อว่าค่ากลางแปรที่ใช้มาตรวัดนี้มีการกระจายแบบปกติก็หมายความว่าสามารถแปลงมาตรวัดแบบลิเคิร์ตให้อยู่ในรูปของคะแนนมาตรฐาน (Z-Scores) ที่มีการกระจายแบบปกติ (Standardized Normal Distribution) ได้ ดังนั้นถ้าระดับความเชื่อมั่น 99.73% สามารถแบ่งค่า Z ออกเป็น 3 ส่วนโดยมีค่ากลางเท่ากับศูนย์ ในขณะที่มาตรวัดลิเคิร์ตแบบห้าตัวเลือกที่สัมพันธ์กับค่า Z-Scores จะมีค่ากลางที่ (3) และอาจแสดงได้ดังรูปต่อไปนี้



รูปที่ 6: คะแนนมาตรฐานของการกระจายแบบปกติเทียบกับคะแนนมาตรวัดลิเคิร์ตแบบห้าตัวเลือก

ลองคิดดูคร่าวๆ แล้ว หากตัวอย่างมีการกระจายแบบปกติ ค่าความแปรปรวนของมาตรวัดลิเคิร์ตแบบห้าตัวเลือกไม่น่าจะเกินครึ่งภายในค่าความคลาดเคลื่อนมาตรฐาน ± 1 แต่ถ้การกระจายไม่เป็นแบบปกติค่านี้อาจเกินหนึ่งได้ จึงไม่มีเหตุผลเลยที่ค่าความแปรปรวนสูงสุดจะเท่ากับ $(0.5)^2$ หรือ 0.25 ตามสูตรของยามานะ ดังนั้นผู้เขียนจึงค่อนข้างกังขาเห็นด้วยหากนักศึกษาใช้มาตรวัดทัศนคติลิเคิร์ตแบบห้าตัวเลือกเป็นส่วนใหญ่ แต่กลับนำสูตรของยามานะตมสมการที่ (6) มาใช้เพียงเพื่อหลีกเลี่ยงการหาค่าความแปรปรวนของประชากรเท่านั้น เหตุที่ไม่เห็นด้วยเพราะ

ข้อสมมติไม่ได้เป็นไปตามที่ยามานะใช้ในการคำนวณสูตรข้างต้น สูตรของยามานะไม่ใช่จะไม่มีค่าความแปรปรวนติดอยู่แต่ยามานะได้แทนค่าความแปรปรวนด้วยตัวเลขที่คิดว่าเหมาะสมกับสถานการณ์การกระจายของสัดส่วนของตัวอย่าง (p) แต่ไม่เหมาะกับการกระจายของค่าเฉลี่ยของตัวอย่าง ($\bar{X}$) ที่เป็นตัวแปรต่อเนื่องเลย หากจะหาค่าความแปรปรวนของกลุ่มตัวอย่างต้องดูด้วยว่ามาตรวัดเป็นแบบใด ควรพยายามหาค่าความแปรปรวนตามมาตรวัดนั้นมาใช้ในการกำหนดขนาดของตัวอย่าง ซึ่งไม่ได้ยากอย่างที่นักศึกษากลัวเลย

นอกจากสูตรของยามานะที่กล่าวข้างต้นแล้ว ยังมีอีกสูตรหนึ่งซึ่งนักศึกษานิยมอ้างถึง ได้แก่

$$n = \left(\frac{1}{4e^2 / Z_{\alpha}^2} \right)$$

ซึ่งก็ไม่ได้ต่างจากสูตรของยามานะเท่าใด เพียงแต่มีที่มาจากประชากรที่มีขนาดใหญ่มากๆ (Infinite Population) แทนที่จะเป็นการสุ่มตัวอย่างจากประชากรที่มีขนาดเล็กลงมา (Finite Population) ที่มาของสูตรนี้จึงมาจากกรณีที่มีมาตรวัดมี 2 ค่าเช่นเดียวกัน และสามารถคำนวณสูตรให้ได้ค่าขนาดของตัวอย่างตามสมการที่ (3) ดังนี้

$$n = \frac{Z_{\alpha}^2 \pi (1 - \pi)}{e^2}$$

หากแทนค่า $\pi = 5$ และ $(1 - \pi) = 0.5$ ตามที่ยามานะเคยใช้ ก็จะได้

$$n = \frac{Z_{\alpha}^2 (0.5)^2}{e^2}$$

หรือ

$$n = \left(\frac{1}{4e^2 / Z_{\alpha}^2} \right) \quad (7)$$

จะหาค่าความแปรปรวนของกลุ่มตัวอย่างก่อนเริ่มทำงานวิจัยได้อย่างไร

หากจะกลับไปตั้งต้นที่สูตรแรก ในกรณีที่มีมาตรวัดเป็นแบบช่วง (Interval Scale) หรือแบบสัดส่วน (Ratio Scale) ซึ่งถือเป็นตัวแปรต่อเนื่อง (Continuous Variable) ในกรณีที่ใช้การสุ่มตัวอย่าง หรือการใช้การเลือกตัวอย่างแบบใช้ความน่าจะเป็น (Probability Sampling) และประชากรมีขนาดใหญ่ จะได้สูตรในการหาขนาดของตัวอย่าง ดังนี้

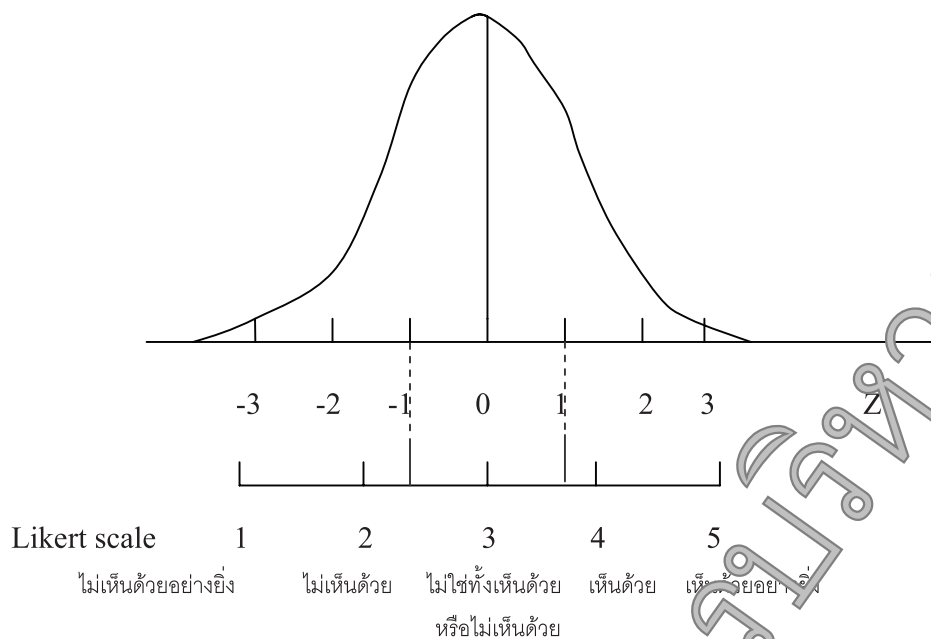
จาก (1)

$$n = \frac{Z_{\alpha}^2 \sigma^2}{e^2} \quad \text{หรือ} \quad n = \frac{Z_{\alpha}^2 s^2}{e^2}$$

เมื่อ n เป็นขนาดของตัวอย่าง
 α เป็นค่า Z-Scores ซึ่งสัมพันธ์กับระดับความเชื่อมั่น $(1 - \alpha)$

σ^2 เป็นค่าความแปรปรวนของประชากร และหากไม่ทราบค่านี้ สามารถใช้ค่าความแปรปรวนของกลุ่มตัวอย่าง (s^2) แทนได้
 e เป็นค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง $(\bar{x} - \mu)$ เมื่อ $\bar{x}$ เป็นค่าเฉลี่ยของกลุ่มตัวอย่าง และ μ เป็นค่าเฉลี่ยของประชากร

เหตุผลที่นักศึกษไทยพบบ่อยคือหลักเลียงที่จะใช้สูตรนี้คงเป็นเพราะไม่ทราบค่าของความแปรปรวนของประชากร (σ^2) หรือของกลุ่มตัวอย่าง (s^2) ก่อนที่จะมีการเลือกตัวอย่างจริงได้อย่างแน่นอน อีกทั้งนั้นก็ไม่ว่าจะกำหนดค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) ให้เป็นเท่าไร ลองมาดูที่ค่า e กัน จะเห็นว่าค่านี้จะขึ้นอยู่กับหน่วยที่ใช้ในการวัดค่าเฉลี่ยของกลุ่มตัวอย่าง ($\bar{x}$) หรือค่าเฉลี่ยของประชากร (μ) ซึ่งก็มีหลากหลายแล้วแต่หน่วยวัดของตัวแปรในแบบที่เราสนใจแต่ละข้อ ดังนั้นการที่จะหาค่า ซึ่งเป็นค่าสมบูรณ์ (Absolute Value) และขึ้นอยู่กับหน่วยวัดของแต่ละตัวแปรก็อาจเริ่มการเก็บข้อมูลจึงค่อนข้างที่จะเป็นปัญหาของนักศึกษส่วนใหญ่ จากประสบการณ์การสอนของผู้เขียนพบว่าในกรณีของนักศึกษปริญญาตรีอาจทำได้ลำบาก เนื่องจากงานวิจัยของนักศึกษปริญญาตรีมักเป็นงานวิจัยเชิงประยุกต์ (Applied Research) ซึ่งแบบสอบถามแต่ละข้อก็มีความวัดที่หลากหลายเกินกว่าที่จะสรุปภาพรวมได้ว่าควรมีมาตรวัดเป็นแบบใด ดังนั้นเพียงแค่จะหาค่า e ก็ค่อนข้างมีปัญหาแล้ว แต่การที่จะหยิบตัวเลขตัวใดตัวหนึ่งมาใส่ก็ไม่ได้ผิดอะไรนัก แต่ต้องยอมให้ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) เมื่อเทียบเป็นร้อยละของค่าเฉลี่ยของกลุ่มตัวอย่างมีขนาดเล็กใหญ่ตามไปด้วย สำหรับงานวิจัยของนักศึกษปริญญาโท ซึ่งส่วนใหญ่เป็นงานวิชาการ (Basic Research) โดยเฉพาะนักศึกษที่ทำงานวิจัยทางสังคมศาสตร์ที่เกี่ยวข้องกับการวัดทัศนคติของกลุ่มเป้าหมาย มักจะต้องวัดตัวแปรที่ไม่สามารถสังเกตได้ (Unobserved Variable หรือ Latent Variable) และมีการใช้มาตรวัดแบบหลายรายการ (Multi-Item Measurement) ซึ่งผู้เขียนพบว่านักศึกษมักใช้มาตรวัดลิเคิร์ตแบบห้าตัวเลือกเป็นส่วนใหญ่ ในกรณีนี้ผู้เขียนมีความเห็นว่าสามารถประมาณการค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) ได้ดังนี้



รูปที่ 7: คะแนนมาตรฐานที่ ± 1 เมื่อแปลงเป็นคะแนนของมาตรวัดลิเคิร์ทแบบห้าตัวเลือก

หากจะพิจารณามาตรวัดลิเคิร์ทแบบห้าตัวเลือก ค่าต่ำสุดที่เป็นไปได้จะเป็น (1) และค่าสูงสุดที่เป็นไปได้จะเท่ากับ (5) หากลองสมมติต่อว่ากลุ่มตัวอย่างมีการกระจายแบบปกติ (Normal Distribution) ค่าถัวเฉลี่ยของกลุ่มตัวอย่างจะเท่ากับ (3) ดังนั้นหากผู้วิจัยยอมให้ความคลาดเคลื่อนจากการสุ่มตัวอย่างมีได้ 5% ของค่าถัวเฉลี่ยจะได้ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) เท่ากับ 5% ของ 3 หรือ 0.15 อย่างไรก็ตามการจะหวังให้การกระจายของกลุ่มตัวอย่างเป็นแบบปกติ (Normal) คงเป็นเรื่องยาก จากประสบการณ์ของผู้เขียนซึ่งเคยใช้ตัวอย่างถึง 2,000 ตัวอย่างในงานวิจัยการกระจายของหลายๆ ตัวแปรที่วัดยังเบ้ซ้ายเบ้ขวาบ้าง ไม่ได้มีการกระจายแบบปกติแต่อย่างใด สำหรับมาตรวัดแบบ 5 จุด ผู้เขียนลองมาวิเคราะห์ดูแล้ว โอกาสที่ค่าถัวเฉลี่ยของกลุ่มตัวอย่างจะเท่ากับ (1) หรือ (5) แทบไม่มีเลย แต่ถ้าจะได้ค่าถัวเฉลี่ยของกลุ่มตัวอย่างมีค่าคั่นไปทาง (2) หรือ (4) ยังมีโอกาสเป็นไปได้มากกว่า ดังนั้นหากคำนวณค่าความคลาดเคลื่อนจากการสุ่มตัวอย่างที่ 5% ของค่าถัวเฉลี่ยได้คือกรณีค่าถัวเฉลี่ยของกลุ่มตัวอย่างเท่ากับ 2 ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) จะเท่ากับ 5% ของ 2 ซึ่งมีค่า 0.10 และในกรณีที่ค่าถัวเฉลี่ยของกลุ่มตัวอย่างเท่ากับ 4 ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) จะเท่ากับ 5% ของ 4 ซึ่งมีค่า 0.20 และเมื่อพิจารณาสูตรในการหาขนาดของตัวอย่างข้างต้นจะพบว่าค่าขนาดของตัวอย่าง (n) จะแปรผกผันกับ ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) ดังนั้น

ถ้า n ยิ่งมีขนาดเล็ก n จะยิ่งมีขนาดใหญ่ หากจะขอยืมแนวคิดอนุรักษ์นิยม (Conservative) ของคุณยามานะที่แทนค่า π และ $(1-\pi)$ ด้วย 0.5 ซึ่งเป็นค่าที่ทำให้ค่าความแปรปรวนของสัดส่วนมีค่าสูงสุดที่ 0.25 ในกรณีนี้ก็น่าจะใช้ค่า $e = .10$ มากกว่าที่จะใช้ค่า e ที่เท่ากับ .15 หรือ .20 เพราะจะทำให้ขนาดของตัวอย่างซึ่งผกผันกับ e มีค่าสูง



ส่วน σ^2 นั้นคงต้องใช้ s^2 ซึ่งได้จากกลุ่มตัวอย่าง แทนที่อยู่แล้ว เพราะปกติก็ไม่น่าทราบค่าความแปรปรวนของ ประชากร (σ^2) อย่างไรก็ตามแม้จะต้องการใช้ค่าความแปรปรวนของกลุ่มตัวอย่าง (s^2) แทน หากยังไม่ได้เริ่มงานวิจัยจะไม่ทราบค่านี้เช่นกัน หากจะย้อนกลับไปดูข้อมูลงานวิจัยเก่าที่ใช้มาตรวัดแบบ 5 จุด ค่า s^2 ที่ได้จากงานวิจัยจะหลากหลายมาก แต่ผู้วิจัยคงพอสรุปค่า s^2 ส่วนใหญ่ที่ได้จากงานวิจัยที่ผ่านมาได้ แม้ว่าการวิจัยที่ต่างกันจะให้ค่า s^2 ที่ต่างกันด้วย หากผู้วิจัยมีการอ้างอิงที่มาของค่า s^2 ที่ใช้ว่าได้มาจากงานของนักวิจัยใด น่าจะใช้ค่า s^2 นั้นเป็นจุดตั้งต้นในการคำนวณขนาดของตัวอย่างที่ยอมรับได้ สำหรับงานวิจัยของผู้วิจัยที่ผ่านมาที่ใช้มาตรวัดแบบนี้จะขึ้นอยู่กับกลุ่มเป้าหมายว่าเป็นใคร ซึ่งพอสรุปได้ว่าค่า s^2 ส่วนใหญ่จะอยู่ระหว่าง 1 ถึง 2 หากจะมองกลุ่มตัวอย่างที่มีการกระจายแบบปกติแล้วจะสามารถแบ่งมาตรวัดลิเคิร์ทแบบห้าตัวเลือกออกได้เป็น 4 ช่วงดังรูปที่ 7 โดยมีค่ากลางอยู่ที่ (3) ดังนั้นทางด้านซ้ายของค่า (3) จะมีอยู่ 2 ช่วง และทางด้านขวาจะมี 2 ช่วง หากจะแปลงมาตรวัดลิเคิร์ทแบบห้าตัวเลือกออกเป็นค่า Z-Scores ที่มีค่ากลางอยู่ที่ 0 และมีค่าความคลาดเคลื่อนมาตรฐาน

(Standard Deviation) ที่ ± 1 ภายใน Z-Scores ที่เท่ากับ ± 3 โอกาสที่ค่าถัวเฉลี่ยของประชากรจะตกอยู่ภายในช่วงนี้ถึง 99.73% ซึ่งหมายความว่าสามารถแบ่งค่า Z ออกเป็นช่วง 3 ช่วง เท่ากับว่าหากหารช่วงของมาตรวัด (1) ถึง (5) ด้วย 6 ค่า Z ที่เท่ากับ +1 หรือ -1 จะเทียบเป็นมาตรวัดลิเคิร์ทแบบห้าตัวเลือกได้ประมาณ 0.67 หน่วย ($1/6 = 0.667$) ซึ่งคำนวณเป็นค่าความแปรปรวนของกลุ่มตัวอย่าง (s^2) ได้ประมาณ 0.45 หน่วยเท่า (นั่นคือ $(0.67)^2 = 0.45$) อย่างไรก็ตามโอกาสที่ผู้วิจัยจะเฝ้าระวังการกระจายของกลุ่มตัวอย่างเป็นแบบปกติเป็นน้อยมาก ดังนั้นผู้เขียนขอใช้วิจารณ์งานที่ได้จากการเฝ้าระวังงานวิจัยที่เคยใช้มาตรวัดลิเคิร์ทแบบห้าตัวเลือกในการประมาณการค่าความแปรปรวนของกลุ่มตัวอย่าง (s^2) น่าจะอยู่ระหว่าง 1 ถึง 2 ดังนั้นหากจะใช้ระดับความเชื่อมั่นที่ 95% จะได้ค่า $Z = \pm 1.96$ และถ้าจะคิดให้ระมัดระวัง (Conservative) คือให้ขนาดของตัวอย่างใหญ่พอที่จะรองรับระดับความเชื่อมั่นที่ต้องการโดยยอมให้ความคลาดเคลื่อนจากการสุ่มตัวอย่างเท่ากับ 5% ของค่าถัวเฉลี่ยของตัวอย่าง ผู้เขียนขอใช้ $e = .10$ ด้วยเหตุผลที่กล่าวมาแล้ว จะสามารถหาขนาดของตัวอย่างได้ดังนี้

จากสูตร	$n = \frac{Z_{\alpha}^2 s^2}{e^2}$	
จะได้	$n = \frac{(1.96)^2 (1)}{(0.10)^2} = 384$	เมื่อ $s^2 = 1$ และ
	$n = \frac{(1.96)^2 (2)}{(0.10)^2} = 768$	เมื่อ $s^2 = 2$

ในกรณีทั้งปวงประมาณในการสุ่มตัวอย่างมีน้อยถ้ายอมให้ความคลาดเคลื่อนจากการสุ่มตัวอย่างเพิ่มเป็น 10% ของค่าถัวเฉลี่ยของกลุ่มตัวอย่าง เมื่อสมมติให้ค่าถัวเฉลี่ยของ

กลุ่มตัวอย่างอยู่ใกล้ๆ 2 จะได้ค่า $e = .20$ ซึ่งจะได้ค่าขนาดของตัวอย่างดังนี้

$n = \frac{(1.96)^2 (1)}{(0.20)^2} = 96$	เมื่อ $s^2 = 1$ และ
$n = \frac{(1.96)^2 (2)}{(0.20)^2} = 192$	เมื่อ $s^2 = 2$

ความจริง หากจะพิจารณาสูตรของยามานะ หากแทนค่า $e = 0.05$ เท่ากับว่ามีการใช้ค่าความคลาดเคลื่อนจากการสุ่มตัวอย่างที่ 10% ของค่าถัวเฉลี่ย [$e/p = 0.05/0.5 = 10\%$]

อันเนื่องมาแต่สูตรของยามาเน

อย่างไรก็ตามในกรณีที่คำนวณขนาดของตัวอย่างได้น้อยกว่า 100 ตัวอย่าง ผู้เขียนมีความเห็นว่าควรปรับขนาดของตัวอย่างให้ได้อย่างน้อย 100 ตัวอย่างสำหรับงานวิจัยเชิงปริมาณ เพราะในกรณีที่ประมวลผลออกมาเป็นค่าร้อยละ ขนาดของตัวอย่างขั้นต่ำไม่ควรต่ำกว่า 100 ตัวอย่าง อันที่จริงมีหลายครั้งที่ผู้เขียนไม่ได้ใช้สูตรในการหาขนาดของตัวอย่างเลย โดยเฉพาะอย่างยิ่งเมื่อมีการเลือกตัวอย่างแบบไม่อาศัยความน่าจะเป็น (Non-Probability Sampling) จากประสบการณ์งานวิจัยที่เคยทำ จากการพูดคุยกับนักวิจัยด้วยกัน และจากการอ่านตำราวิจัยตลาดภาษาอังกฤษหลายๆ เล่ม หากไม่มีความจำเป็นต้องวิเคราะห์กลุ่มย่อย พอสรุปได้ว่า ควรเลือกตัวอย่างประมาณ 200 ตัวอย่างเป็นอย่างน้อย อย่างไรก็ตามเมื่อมีการเก็บตัวอย่างของจริงเสร็จแล้ว การคำนวณค่าความแปรปรวนสูงสุดของกลุ่มตัวอย่าง (s^2) ที่ได้จากทุกคำถามในแต่ละงานวิจัยจะทำได้ง่ายอยู่แล้ว แม้ว่าค่าความแปรปรวนของกลุ่มตัวอย่าง (s^2) นั้นจะแตกต่างจากสิ่งที่ผู้วิจัยคาดเดาไว้ตั้งแต่ตอนจัดทำข้อเสนองานวิจัย สิ่งที่ผู้วิจัยต้องยอมเสียก็คือค่าความคลาดเคลื่อนจากการสุ่มตัวอย่าง (e) ที่ต่างไปจากค่าที่ระบุไว้ตอนเขียนข้อเสนอ

โครงการวิจัย ทั้งนี้ค่านี้อาจเพิ่มขึ้นหรือลดลงได้เนื่องจากค่าความแปรปรวนของกลุ่มตัวอย่าง (s^2) ที่สามารถหาได้จริงอาจสูงหรือต่ำกว่าค่าที่ใช้ในการคำนวณขนาดของตัวอย่างตามสูตรที่ใช้ในช่วงก่อนการเก็บตัวอย่าง แต่ผู้เขียนยังไม่เคยเห็นผู้อ่านท่านใดตามไปตรวจสอบจุดนี้เลย

ในการกำหนดขนาดของตัวอย่าง นอกจากการใช้สูตรหรือการใช้วิจารณ์กำหนดขนาดของตัวอย่างตามประสบการณ์ที่มีแล้ว ประเภทของสถิติที่ใช้ในการวิเคราะห์ข้อมูลก็มีส่วนในการกำหนดขนาดของตัวอย่างเช่นกัน ยกตัวอย่างเช่น ในกรณีที่ผู้วิจัยใช้สถิติการวิเคราะห์องค์ประกอบ (Exploratory Factor Analysis) มีผู้ระบุว่า “หลักปฏิบัติคร่าวๆ คือ ควรใช้ขนาดของตัวอย่างอย่างน้อย 4-5 เท่าของจำนวนตัวแปรที่ใช้ในการวิเคราะห์” (Basilesky, 1994; อ้างถึงใน Malhotra, 2010, p.639) นั่นก็หมายความว่า ถ้าต้องใช้แบบแผนการวิเคราะห์องค์ประกอบ 30 ตัวแปร ขนาดของตัวอย่างที่เก็บไม่ควรน้อยกว่า 150 ตัวอย่างโดยประมาณนั้นเอง

เอกสารอ้างอิง

Lindgren, Bernard W. *Statistical Theory*. 3rd ed. New York: Macmillan Publishing Co., Inc., 1976.

Malhotra, Naresh K. *Marketing Research: An Applied Approach*.

Orientation. 6th ed. Upper Saddle River: Pearson Education, Inc., 2010.

Yamane, Taro. *Statistics: An Introductory Analysis*. 3rd ed. New York: Harper &

Row, Publishers, Inc., 1973.

Zinkmund, William G. *Exploring Marketing Research*. 8th ed. Mason: South-Western, 2003.