

อีกหนึ่งทางเลือกที่ทำให้การเรียนรู้สถิติอนุมาณน่าสนใจมากขึ้น

An Alternative Way to Make the Learning of Inferential Statistics More Interesting

รมิดา ศรีเหรา*

ภาควิชาคณิตศาสตร์และสถิติ คณะวิทยาศาสตร์และเทคโนโลยี

มหาวิทยาลัยธรรมศาสตร์ ศูนย์รังสิต ตำบลคลองหนึ่ง อำเภอคลองหลวง จังหวัดปทุมธานี 12120

Ramidha Srihera*

Department of Mathematics and Statistics, Faculty of Science and Technology,

Thammasat University, Rangsit Centre, Klong Nueng, Klong Luang, Pathum Thani 12120

บทคัดย่อ

การสอนเรื่องสถิติพรรณนาและสถิติอนุมานมีความแตกต่างกัน สถิติพรรณนาเป็นเรื่องที่ทำความเข้าใจได้ง่ายกว่า สำหรับสถิติอนุมานนักศึกษามักจะคิดว่ามีสัญลักษณ์และสูตรสถิติต่าง ๆ ที่ยากและสับสน ด้วยเหตุนี้ทำให้พวกเขาเรียนรู้ด้วยการท่องจำโดยปราศจากความเข้าใจที่แท้จริง ทั้งนี้เนื่องจากการขาดความเชื่อมโยงระหว่างแนวคิดพื้นฐานที่สำคัญกับเนื้อหาในบทเรียน ดังนั้นบทความนี้ได้กล่าวถึงแนวคิดพื้นฐานทางสถิติในเชิงทฤษฎีที่สำคัญ 5 ข้อ ได้แก่ การชักตัวอย่างแบบสุ่ม ความคลาดเคลื่อนจากการสุ่ม การแจกแจงของสิ่งตัวอย่างของค่าเฉลี่ยตัวอย่าง การแจกแจงแบบปกติและทฤษฎีขีดจำกัดส่วนกลาง ซึ่งสาระสำคัญอยู่ที่การนำแนวคิดดังกล่าวมาอธิบายกับเรื่องการประมาณค่าเฉลี่ยของประชากรโดยใช้ตัวอย่างจากงานวิจัย ทั้งนี้เพื่อให้ นักศึกษาสามารถเข้าใจในส่วนสำคัญส่วนหนึ่งของสถิติอนุมานได้ดียิ่งขึ้น

คำสำคัญ : การชักตัวอย่างแบบสุ่ม, ความคลาดเคลื่อนจากการสุ่ม, การแจกแจงของสิ่งตัวอย่างของค่าเฉลี่ยตัวอย่าง, การแจกแจงแบบปกติ, ทฤษฎีขีดจำกัดส่วนกลาง การประมาณค่าเฉลี่ยของประชากร

Abstract

Teaching descriptive and inferential statistics are different. It seems that descriptive statistics is much easier to understand. According to inferential statistics, what students have in common is statistical symbols and formulas are unbearable difficult and confusing. Therefore, they learn statistics by remembering without true understanding. It is due to lack of the connection between statistical concepts and lesson content. Consequently, this article presents mainly in five elementary statistical concepts such as random sampling, sampling error, the sampling distribution of the sample mean, the normal distribution and the central limit theorem. Moreover, taking these concepts to

describe estimating the population mean from research work is highlighted so that students are able to obtain a better understand of one of the most important part in inferential statistics.

Keywords: random sampling, sampling error, sampling distribution of the sample mean, normal distribution, central limit theorem, estimating the population mean

1. บทนำ

โดยทั่วไปเอกสารทางวิชาการ (ตำราหรือหนังสือ) วิชาสถิติพื้นฐานไม่ว่าจะถูกเขียนขึ้นสำหรับนักศึกษาของสาขาวิชาใดก็ตาม การเรียงลำดับของบทเรียนจะเริ่มต้นด้วยความรู้เบื้องต้นเกี่ยวกับสถิติซึ่งจะครอบคลุมถึงสถิติพรรณนา ตามมาด้วยเรื่องของความน่าจะเป็น ตัวแปรสุ่มและการแจกแจงแบบปกติ การชักตัวอย่างและการแจกแจงของค่าที่ได้จากตัวอย่าง แล้วจึงเข้าสู่บทเรียนของสถิติอนุมาน ซึ่งประกอบด้วย การประมาณค่าและการทดสอบสมมติฐานของค่าพารามิเตอร์ ความแตกต่างของเอกสารเหล่านั้นจะอยู่ที่ความเข้มข้นของเนื้อหาทางทฤษฎี สังเกตได้ว่าหากเป็นเอกสารทางวิชาการที่ใช้สำหรับสอนนักศึกษาสาขาวิชาสถิติจะเน้นหนักไปทางทฤษฎี การยกตัวอย่างประกอบมีไม่มาก ในทางตรงกันข้ามหากเป็นเอกสารวิชาการที่ใช้สำหรับสอนนักศึกษาในสาขาวิชาอื่น อาทิเช่น สังคมศาสตร์ วิศวกรรมศาสตร์ หรือแพทยศาสตร์ จะเน้นที่ตัวอย่างมากมายที่เกี่ยวข้องกับสาขาวิชาที่เรียน มีการกล่าวถึงทฤษฎีไม่มาก ครั้นเมื่อเข้าสู่บทเรียนของสถิติอนุมาน (มักจะเริ่มเรียนประมาณกลางภาคการศึกษา) นักศึกษาจะมีความรู้ลึกกว่าบทเรียนนี้เข้าใจอย่างมาก มีสัญลักษณ์และสูตรต่าง ๆ จำนวนมากที่ยากและสับสน (Salkind, 2008) ความยากที่กล่าวถึงนั้นแตกต่างกันตามธรรมชาติของนักศึกษาแต่ละสาขาวิชา นักศึกษาที่เรียนในสาขาวิชาสถิติจะรู้สึกเริ่มเหนื่อยล้ากับทฤษฎีต่าง ๆ ที่เรียนมา ไม่อยากที่จะจดจำหรือทำความเข้าใจในบทเรียนต่อ ๆ

อีกต่อไป การเรียนมีลักษณะเหมือนเรียนแยกส่วน จบในแต่ละบท ได้ผลลัพธ์มาแล้วแปลความหมายไม่ได้ เนื่องจากไม่ได้เห็นตัวอย่างมาก ส่วนนักศึกษาที่เรียนในสาขาวิชาอื่น ๆ เมื่อเรียนในบทเรียนนี้พวกเขาจะรู้สึกเหมือนว่าโดนถล่มด้วยสูตรมากมาย ไม่ค่อยเข้าใจที่มาที่ไปของสูตรต่าง ๆ เพราะไม่ได้เรียนทฤษฎีมากมายพอ กล่าวได้ว่านักศึกษาจากสาขาวิชาใดก็ตามเมื่อเข้าสู่บทเรียนสถิติอนุมานตกอยู่ในสถานการณ์เดียวกันคือมีสิ่งพวกเขาจะมีคำถามมากมายที่เกิดขึ้นขณะเรียนไม่แตกต่างกัน ยกตัวอย่าง เช่น ในการประมาณค่าเฉลี่ยของประชากรแบบช่วงที่ต้องมีข้อตกลงเบื้องต้นว่าตัวอย่างสุ่มจะต้องมีการแจกแจงแบบปกติและใช้การคำนวณมาตรฐานจากตารางการแจกแจงแบบปกติมาตรฐาน แต่หากตัวอย่างสุ่มมีการแจกแจงที่ไม่ใช่การแจกแจงแบบปกติ ทำไม่จึงยังสามารถใช้ค่าจากตารางการแจกแจงแบบปกติมาตรฐานได้อยู่ ซึ่งเมื่อได้รับคำตอบแล้วก็ยังรู้สึกว่าจะไม่เข้าใจอยู่ดี ด้วยเหตุนี้นักศึกษาจึงเลือกที่จะเรียนสถิติอนุมานด้วยการพยายามจำสูตร แทนค่าตัวเลขลงในสูตรต่าง ๆ เพื่อหาคำตอบเพียงอย่างเดียวเท่านั้น ซึ่งแท้จริงแล้วการเรียนโดยใช้ความรู้จากบทเรียนต่าง ๆ ก่อนหน้านั้นจะช่วยให้นักศึกษามีความเข้าใจในเรื่องสถิติอนุมานมากขึ้น

หนังสือเรียนเกี่ยวกับการวิเคราะห์ทางสถิติสำหรับนักศึกษาทางกลุ่มสาขาวิชาทางสังคมศาสตร์หรือด้านวิศวกรรมศาสตร์ที่ออกมาใหม่ ๆ โดยมากพยายามที่จะเขียนเชื่อมโยงระหว่างแนวคิดพื้นฐานทางสถิติที่สำคัญ ได้แก่ การชักตัวอย่างแบบสุ่ม

(random sampling) ความคลาดเคลื่อนจากการสุ่ม (sampling error) การแจกแจงของสิ่งตัวอย่างของค่าเฉลี่ยตัวอย่าง (sampling distribution of the sample mean) การแจกแจงแบบปกติ (normal distribution) และทฤษฎีขีดจำกัดส่วนกลาง (central limit theorem) ในบทเรียนของสถิติอนุมานโดยไม่เน้นทฤษฎี (Cadwell, 2010; Devore, 2012; Howell, 2012) ซึ่งเป็นแนวทางการถ่ายทอดความรู้ที่น่าสนใจ อย่างไรก็ตาม หากได้เพิ่มเติมเนื้อหาในส่วนของทฤษฎีในระดับที่พอเหมาะน่าจะเป็นประโยชน์ภาพในการทำความเข้าใจในบทเรียนสถิติอนุมาน ด้วยเหตุนี้ผู้เขียนจึงพยายามหาจุดสมดุลระหว่างช่องว่างของการผลิตเอกสารทางวิชาการที่ได้กล่าวข้างต้น โดยเลือกที่จะอธิบายความเชื่อมโยงของทฤษฎีทางสถิติเข้ากับตัวอย่างงานวิจัยในการประมาณค่าแบบช่วงของค่าเฉลี่ยของประชากร ทั้งนี้เพื่อให้การถ่ายทอดความรู้ในลักษณะนี้เป็นอีกทางเลือกหนึ่งที่จะทำให้ นักศึกษารู้สึกว่าการเรียนสถิติอนุมานมีความน่าสนใจมากขึ้น

2. การอธิบายแนวความคิดพื้นฐานทางสถิติทั้ง 5 ประการ โดยใช้ตัวอย่างงานวิจัย

สมมติว่า “ต้องการทราบอายุเฉลี่ยของกลุ่มผู้ใช้อินเทอร์เน็ตในประเทศไทย” ในทางปฏิบัติจะต้องทำการสำรวจโดยการสอบถามอายุของผู้ใช้อินเทอร์เน็ตทุกคนที่อาศัยอยู่ในประเทศไทย ซึ่งทางสถิติถือเป็นประชากร (population) ค่าอายุเฉลี่ยของประชากรคือค่าพารามิเตอร์ (parameter) แต่ในความเป็นจริงด้วยข้อจำกัดต่าง ๆ อาทิเช่น เวลา งบประมาณ ไม่สามารถที่จะสอบถามประชากรทุกคนได้ สิ่งที่สามารถทำได้คือการประมาณค่าพารามิเตอร์ (parameter estimation) ซึ่งกรณีนี้ต้องการประมาณอายุเฉลี่ยของกลุ่มผู้ใช้อินเทอร์เน็ตในประเทศไทยทุกคน (estimating the

population mean) จากผู้ใช้อินเทอร์เน็ตบางคนที่อาศัยอยู่ในประเทศไทย คือ ตัวอย่าง (sample) ค่าอายุเฉลี่ยของตัวอย่าง คือ ค่าสถิติ (statistic) จะถูกคำนวณขึ้นเพื่อประมาณค่าพารามิเตอร์

จากรายงานผลการสำรวจกลุ่มผู้ใช้อินเทอร์เน็ตในประเทศไทยปี 2553 ได้สุ่มสอบถามผู้ใช้อินเทอร์เน็ตจำนวน 14,067 คน (ศูนย์เทคโนโลยีอิเล็กทรอนิกส์และคอมพิวเตอร์แห่งชาติ, 2553) หากต้องการให้สิ่งที่ประมาณใกล้เคียงกับความเป็นจริงมากที่สุด ตัวอย่างที่ใช้ควรจะเป็นตัวอย่างที่ได้จากการชัก ตัวอย่างแบบสุ่ม (random sampling) และตัวอย่างที่ได้โดยวิธีนี้ เรียกว่า ตัวอย่างสุ่ม (random sample)

นิยาม : ตัวอย่างสุ่ม (random sample) ขนาด n จากประชากรที่มีความหนาแน่น $f(x)$ หมายถึงเซตของตัวแปรสุ่ม $x_1, x_2, \dots, x_n$ ที่เป็นอิสระกันและแต่ละตัวแปรสุ่มมีความหนาแน่นเช่นเดียวกับความหนาแน่นของประชากร (ประชุม, 2545) จากนิยามสามารถอธิบายตัวอย่างสุ่มให้เข้าใจง่ายขึ้น ดังนี้ (1) ผู้ใช้อินเทอร์เน็ตในประเทศไทยทุกคนมีโอกาสที่จะถูกเลือกเป็นตัวอย่าง (2) การสอบถามผู้ใช้อินเทอร์เน็ตเป็นไปอย่างสุ่ม เช่น สมมติว่าผู้ถูกเลือกเป็นตัวอย่างคนแรกเลือกจากผู้ที่ใช้ในกรุงเทพมหานคร ตัวอย่างคนที่สองที่จะถูกเลือกโดยไม่ได้รับอิทธิพล (เป็นอิสระ) จากการเลือกตัวอย่างคนแรก กล่าวคือ คนที่สองอาจถูกเลือกจากกรุงเทพฯ หรือจากจังหวัดใดในประเทศไทยก็ได้ ข้อมูลที่สอบถามคืออายุ (x) ดังนั้นอายุจะเป็นตัวแปรสุ่มซึ่งมีทั้งสิ้น $x_1, x_2, \dots, x_{14,067}$ ตัวซึ่งถือเป็นหนึ่งตัวอย่างสุ่ม การใช้ตัวอย่างสุ่มทำให้สามารถวิเคราะห์ข้อมูลด้วยวิชาการสถิติ ทฤษฎีต่าง ๆ ทางสถิติได้ต่อไป (ประชุม, 2552) ถึงแม้ว่าในงานวิจัยแต่ละชิ้นจะใช้เพียงหนึ่งตัวอย่างสุ่ม แต่ในความเป็นจริงแล้วมีเซตของตัวอย่างสุ่มขนาด n มากกว่าหนึ่ง

เซต หากมีการชักตัวอย่างแบบสุ่มมากกว่าหนึ่งครั้ง จำนวนของค่าสถิติที่คำนวณได้จะมีค่าเท่ากับจำนวนเซตของตัวอย่างสุ่ม ซึ่งหากสมมติว่าทราบค่าพารามิเตอร์ จะสามารถทราบได้ว่าค่าสถิติที่คำนวณมีความแตกต่างจากค่าพารามิเตอร์หรือไม่

นิยาม : ความคลาดเคลื่อนจากการสุ่ม (sampling error) หมายถึง ค่าความแตกต่างระหว่างค่าสถิติและค่าพารามิเตอร์

พิจารณาจากงานวิจัย ถ้าทำการสุ่มถามผู้ใช้ อินเทอร์เน็ตในประเทศไทย 14,067 คน ซ้ำ ๆ กันหลายครั้ง และในการสุ่มแต่ละครั้งจะคำนวณค่าอายุเฉลี่ยของตัวอย่าง (ค่าสถิติ) ใ้ทุกครั้ง ตัวอย่างเช่น ถ้าสุ่มตัวอย่างจำนวน 5 ครั้ง ค่าความอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยที่ถูกเลือกเป็นตัวอย่างแต่ละครั้งได้เท่ากับ 32.8, 35.5, 25.3, 18.7 และ 33.4 ปี สมมติทราบว่าค่าอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยทั้งหมด (ค่าพารามิเตอร์) เท่ากับ 33.4 ปี จะเห็นมีความเป็นไปได้ที่ค่าสถิติมีค่าต่ำกว่า สูงกว่า หรือแม้กระทั่งเท่ากับค่าพารามิเตอร์ หากเมื่อใดที่ค่าสถิติมีค่าไม่เท่ากับค่าพารามิเตอร์ นั้นแสดงว่าความคลาดเคลื่อนจากการสุ่มได้เกิดขึ้นแล้ว

และหากลองขยายจำนวนการชักตัวอย่างจาก 5 ครั้ง เป็น 5,000 ครั้ง จะพบว่าตัวเลขค่าอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยในที่คำนวณจากแต่ละตัวอย่างสุ่มจะเพิ่มขึ้นจาก 5 ค่า เป็น 5,000 ค่า ซึ่งจำนวนตัวเลขมากมายนี้เป็นประโยชน์อย่างมากเพราะทำให้ทราบลักษณะธรรมชาติของข้อมูลที่กำลังศึกษาอยู่

นิยาม : การแจกแจงของสิ่งตัวอย่างของค่าเฉลี่ยตัวอย่าง (sampling distribution of the sample mean) หมายถึง การแจกแจงของความน่าจะเป็นของค่าเฉลี่ยตัวอย่างซึ่งคำนวณจากตัวอย่างสุ่มขนาดเดียวกันที่เป็นไปได้ทั้งหมด

เมื่อนำนิยามมาประยุกต์กับงานวิจัยสามารถอธิบายได้ดังนี้ ในการสุ่มซ้ำ 5,000 ครั้ง เพื่อถามผู้ใช้อินเทอร์เน็ตในประเทศไทย 14,067 คน พบว่ามีตัวเลขค่าอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยในแต่ละตัวอย่างสุ่มจำนวน 5,000 ค่า นำตัวเลขทั้งหมดนี้มาสร้างเป็นตารางการแจกแจงความถี่ของค่าอายุเฉลี่ยของตัวอย่างแต่ละค่า หลังจากนั้นนำมาสร้างเป็นแผนภูมิแท่งโดยให้แกน Y เป็นค่าเฉลี่ยของตัวอย่างและแกน X เป็นความถี่ หลังจากนั้นปรับแผนภูมิแท่งเป็นฮิสโตแกรม ต่อจากนั้นปรับฮิสโตแกรมเป็นรูปหลายเหลี่ยมความถี่จนกระทั่งท้ายที่สุดปรับรูปหลายเหลี่ยมความถี่เป็นเส้นโค้งของการแจกแจงของค่าเฉลี่ยของตัวอย่าง ซึ่งการทำเช่นนี้จะทำให้เห็นภาพการกระจายของค่าเฉลี่ยของตัวอย่าง เส้นโค้งดังกล่าวจะมีรูปร่างแตกต่างกันไปตามธรรมชาติของข้อมูล แต่รูปของเส้นโค้งที่มักจะสะท้อนความเป็นธรรมชาติของข้อมูลมากที่สุด (ค่าที่สูงเกินไปหรือต่ำเกินไปมีจำนวนน้อยในสัดส่วนที่เท่า ๆ กัน ส่วนค่าที่อยู่ตรงกลางมีจำนวนมาก) จะเป็นเส้นโค้งที่เป็นรูประฆังคว่ำและมีลักษณะสมมาตร

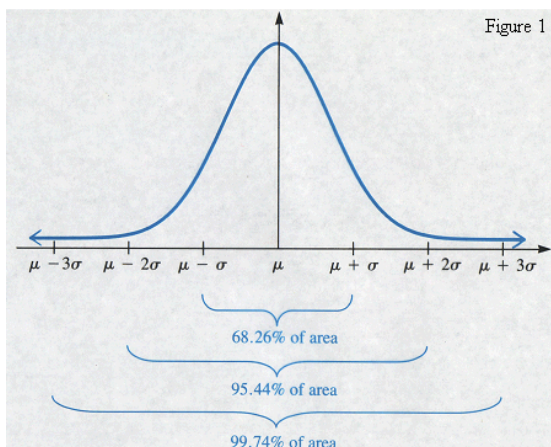
นิยาม : การแจกแจงแบบปกติ (normal distribution) หมายถึง การแจกแจงของความน่าจะเป็นของค่าของตัวแปรสุ่มที่เป็นค่าแบบต่อเนื่อง

โดยฟังก์ชันความหนาแน่นของการแจกแจงแบบปกติ คือ $f(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2\sigma^2}(x-\mu)^2}$; $-\infty < x < \infty, -\infty < \mu < \infty, \sigma^2 > 0$ โดย x แทนตัวแปรสุ่ม

μ แทนค่าเฉลี่ย และ σ^2 แทนค่าความแปรปรวน ซึ่งทั้งสองค่าเป็นพารามิเตอร์

ซึ่งกราฟแสดงค่าฟังก์ชันความหนาแน่น (probability density function) มีลักษณะดังรูปที่ 1 ค่า σ^2 เป็นค่าที่ใช้วัดการกระจายสามารถทำได้โดย

คำนวณหาผลต่างของค่าเฉลี่ยของตัวอย่างแต่ละค่า และค่าเฉลี่ยของค่าเฉลี่ยตัวอย่างแต่ละค่า



รูปที่ 1 เส้นโค้งการแจกแจงแบบปกติที่มีค่าเฉลี่ย μ และความแปรปรวน σ^2

สำหรับการแจกแจงแบบปกติมีค่า $\mu = 0$ และ $\sigma^2 = 1$ จะเรียกว่าการแจกแจงแบบปกติมาตรฐาน (z) (วิกิพีเดีย : สารานุกรมเสรี, 2554)

การทราบว่าตัวอย่างสุ่มที่ศึกษา มีการแจกแจงแบบปกติจะทำให้เราสามารถนำค่าสถิติที่คำนวณได้มาทำการประมาณค่าแบบช่วงและทดสอบสมมติฐานค่าพารามิเตอร์ของประชากรได้อย่างไรก็ตามหากการแจกแจงของสิ่งตัวอย่างของค่าเฉลี่ยตัวอย่างไม่ได้มีการแจกแจงแบบปกติ การประมาณค่าหรือทดสอบสมมติฐานของค่าพารามิเตอร์ซึ่งต้องใช้ตารางการแจกแจงแบบปกติเป็นเรื่องที่ไม่สามารถทำได้ และเป็นเรื่องที่ยุ่งยากมาก ดังนั้น หากต้องการใช้ตัวอย่างสุ่มที่ไม่ได้มีการแจกแจงแบบปกติในการประมาณค่าแบบช่วงหรือทดสอบสมมติฐานของค่าพารามิเตอร์ทำได้โดยอาศัยทฤษฎีที่มีความสำคัญเป็นอย่างมากซึ่งจะได้กล่าวถึงต่อไป

นิยาม : ทฤษฎีขีดจำกัดส่วนกลาง (central limit theorem)

ถ้า $\bar{X}$ เป็นค่าเฉลี่ยของตัวอย่างสุ่ม $X_1, X_2, \dots, X_n$ มีการแจกแจงของความน่าจะเป็นเหมือนกัน โดยมีค่าเฉลี่ย μ และค่าความแปรปรวน σ^2 ดังนั้น $\frac{(\bar{X} - \mu)}{\sigma/\sqrt{n}}$ จะมีการแจกแจงของความน่าจะเป็นแบบปกติมาตรฐาน (z) เมื่อ n มีค่ามาก (Hogg and Tanis, 2010)

กล่าวคือ ในกรณีที่ประชากรมีการแจกแจงของความน่าจะเป็นแบบใดก็ตาม หากขนาดตัวอย่างมีจำนวนมากพอ การแจกแจงของสิ่งตัวอย่างของค่าเฉลี่ยของตัวอย่าง ($\bar{X}$) จะมีการแจกแจงของความน่าจะเป็นเข้าใกล้การแจกแจงของความน่าจะเป็นแบบปกติ ซึ่งค่าเฉลี่ยและค่าส่วนเบี่ยงเบนมาตรฐานของค่าเฉลี่ยของตัวอย่าง ($\mu_{\bar{X}}$ และ $\sigma_{\bar{X}}$) มีค่าเท่ากับค่าเฉลี่ยของประชากร (μ) และค่าส่วนเบี่ยงเบนมาตรฐานของประชากรหารด้วยรากที่สองของขนาดตัวอย่าง ($\sigma/\sqrt{n}$) เรียกว่า ค่าความคลาดเคลื่อนมาตรฐานของค่าเฉลี่ย (standard error of the mean)

3. การเชื่อมโยงแนวคิดพื้นฐานทางสถิติ 5 ข้อ กับการประมาณค่าเฉลี่ยประชากรแบบช่วง (กรณีทราบค่า σ)

คราวนี้ย้อนกลับไปที่ตัวอย่างที่ข้างต้น สมมติว่าคำนวณอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทย จากตัวอย่างสุ่มได้เท่ากับ 25.3 ปี โดยหลักการของสถิติอนุมาน กล่าวได้ว่า ผลจากการสำรวจพบว่า อายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยเท่ากับ 25.3 ปี ซึ่งแท้จริงมีความเป็นไปได้ว่าตัวเลข 25.3 อาจจะเท่ากับ มากกว่า หรือน้อยกว่า อายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยทั้งหมดก็เป็นได้ ถ้าเช่นนั้นจะแน่ใจหรือมั่นใจได้อย่างไรว่าการอ้างอิงผลการสำรวจด้วยตัวเลขเพียงตัวเดียวจะเข้าใกล้หรือห่างไกลจากค่าที่

แท้จริง (Cohen, 2008) นอกจากนี่สิ่งที่จะละเลยไม่ได้คือ ในการประมาณค่าพารามิเตอร์ใด ๆ ก็ตามย่อมมีโอกาสที่จะเกิดความคลาดเคลื่อนจากการสุ่มขึ้นได้เสมอ ดังนั้นจะเป็นการดีกว่าถ้าสามารถกล่าวถึงสิ่งที่ต้องการประมาณเป็นช่วงของตัวเลขด้วยความมั่นใจระดับหนึ่ง คำถามคือ อายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยน่าจะมีค่าอยู่ระหว่างกี่ปีถึงกี่ปี ?

หากแนวคิดการใช้ $\bar{x}$ เพื่อประมาณ μ เป็นแนวคิดที่สมเหตุสมผล การขยายการประมาณค่าของ μ เป็นแบบช่วง ควรใช้ $\bar{x}$ เป็นจุดเริ่มต้นประกอบกับความรู้อันฐานที่กล่าวว่าหากตัวแปรสุ่มมีการแจกแจงแบบปกติค่าเฉลี่ยของตัวแปรสุ่มจะมีการแจกแจงแบบปกติด้วย หรือหากตัวแปรสุ่มไม่ได้มีการแจกแจงแบบปกติ สามารถใช้ทฤษฎีขีดจำกัดส่วนกลางในการหาข้อสรุปว่าการแจกแจงของ $\bar{x}$ จะมีการแจกแจงแบบปกติที่มีลักษณะเป็นโค้งสมมาตร

นั่นคือ ค่าประมาณแบบช่วงของ μ ควรจะมีค่าอยู่ระหว่าง $\bar{x} \pm$ ค่าบางค่า ? และค่าเฉลี่ยของประชากร (μ) ควรจะตกอยู่ระหว่างช่วงที่เรากำลังจะสร้างขึ้นด้วยความมั่นใจหรือความเชื่อมั่นเท่าใด ?

หากพิจารณารูปเส้นโค้งการแจกแจงแบบปกติข้างต้น กล่าวได้ว่า หากทำการสุ่มตัวอย่างสุ่มซ้ำ 100 ครั้ง จะมีประมาณ 68 ครั้ง ที่ตัวแปรสุ่ม x จะมีค่าอยู่ระหว่าง $\mu \pm 1\sigma$ ซึ่งตัวเลข 68 เป็นตัวเลขที่บอกถึงความเชื่อมั่นในการประมาณค่าตัวแปรสุ่ม x เราใช้คำว่าด้วยความเชื่อมั่น 68 % ซึ่งสามารถอธิบายได้เช่นเดียวกันสำหรับระดับความเชื่อมั่นที่ 95 % และ 99 % ที่ตัวแปรสุ่ม x จะมีค่าอยู่ระหว่าง $\mu \pm 2\sigma$

และ $\mu \pm 3\sigma$ ตามลำดับ ตัวเลข 1, 2 และ 3 ได้มาจากการปัดเศษทศนิยมของค่าที่ได้จากการเปิดตารางการแจกแจงแบบปกติมาตรฐาน (z) และพบว่า นักสถิติส่วนมากมักจะใช้ระดับความเชื่อมั่น

95 % และ 99 % ในการประมาณค่าแบบช่วง (Pyrczak, 1995)

ดังนั้น ค่าประมาณแบบช่วงของ μ ควรจะอยู่ระหว่าง $\bar{X} \pm Z\sigma_{\bar{x}}$

จากทฤษฎีขีดจำกัดส่วนกลาง (central limit theorem) ค่าส่วนเบี่ยงเบนมาตรฐานของค่าเฉลี่ยของตัวอย่าง ($\sigma_{\bar{x}}$) มีค่าเท่ากับ ค่าความคลาดเคลื่อนมาตรฐานของค่าเฉลี่ย (standard error of the mean) $(\sigma/\sqrt{n})$ ซึ่งหากขนาดตัวอย่างมีจำนวนมากจะส่งผลให้ค่าความคลาดเคลื่อนมาตรฐานของค่าเฉลี่ยลดลงด้วย

ดังนั้น สมมติว่าอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยที่ได้จากตัวอย่างสุ่ม เท่ากับ 32.8 ปี และทราบว่าอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยมีการแจกแจงแบบปกติด้วยค่าเฉลี่ย (μ) เท่ากับ 33.4 ปีและ ส่วนเบี่ยงเบนมาตรฐาน (σ) เท่ากับ 157.64 ปี

ค่าประมาณแบบช่วงของอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยที่ระดับความเชื่อมั่น 95 % จะมีค่าอยู่ในช่วง $2.8 \pm (2)(157.64/\sqrt{14,067}) = (30.14, 35.46)$ ปี หมายความว่า หากใช้ตัวอย่างสุ่มจำนวน 100 ชุด พบว่า 95 ครั้ง ของการประมาณอายุเฉลี่ยของผู้ใช้อินเทอร์เน็ตในประเทศไทยจะมีค่าอยู่ระหว่าง 30.14 ถึง 35.46 ปี

4. บทสรุป

การสอนเรื่องสถิติอนุมานเป็นเรื่องที่ทำหยาบความนี้ได้นำเสนอเฉพาะส่วนของการนำแนวคิดพื้นฐานทางสถิติ 5 ข้อ มาประยุกต์ใช้กับการประมาณค่าเฉลี่ยแบบช่วง (กรณีทราบค่า σ) เทคนิคสำคัญที่ใช้ในการนำเสนอคือความพยายามรักษาสมดุลของความหนักเบาของเนื้อหาและการเชื่อมโยงแนวคิด ทฤษฎีทางสถิติ และตัวอย่างงานวิจัยเข้าด้วยกันในสัดส่วนที่พอดี นอกจากนั้นการ

เรียงลำดับและการผสมผสานเนื้อหาจะช่วยลดความ
ความหนักของเนื้อหาทางทฤษฎี ซึ่งถ้าหาก
นักศึกษามีความเข้าใจในส่วนสิ่งที่อธิบายข้างต้น
ได้ดี ความสามารถในการต่อยอดหรือขยายความรู้
ไปสู่ความเข้าใจในเรื่องของการประมาณค่าเฉลี่ย
แบบช่วง (กรณีไม่ทราบค่า σ) และการทดสอบ
สมมติฐานของค่าพารามิเตอร์ซึ่งเป็นเนื้อหาในสถิติ
อนุมานจะยังมีประสิทธิภาพมากขึ้น

5. กิตติกรรมประกาศ

ผู้เขียนได้รับคำแนะนำและข้อเสนอแนะซึ่ง
เป็นประโยชน์อย่างยิ่งในการเขียนบทความนี้จาก
ผู้ทรงคุณวุฒิผู้อ่านบทความ ซึ่งผู้เขียนขอขอบ
พระคุณมา ณ ที่นี้ด้วย

6. เอกสารอ้างอิง

ประชุม สุวดี, 2545, ทฤษฎีการอนุมานเชิงสถิติ,
โครงการส่งเสริมเอกสารวิชาการ สถาบัน
บัณฑิตพัฒนบริหารศาสตร์, กรุงเทพฯ.
ประชุม สุวดี, 2552, การสำรวจด้วยตัวอย่าง การ
ชักตัวอย่าง และการวิเคราะห์, โครงการ
ส่งเสริมเอกสารวิชาการ สถาบันบัณฑิตพัฒน
บริหารศาสตร์, กรุงเทพฯ.
วิกิพีเดีย : สารานุกรมเสรี, การแจกแจงแบบปกติ,
แหล่งที่มา : <http://th.wikipedia.org/wiki>, 16
กันยายน 2554.

ศูนย์เทคโนโลยีอิเล็กทรอนิกส์และคอมพิวเตอร์
แห่งชาติ, 2553, รายงานผลการสำรวจกลุ่ม
ผู้ใช้อินเทอร์เน็ตในประเทศไทยปี 2553,
สำนักงานพัฒนาวิทยาศาสตร์และเทคโนโลยี
แห่งชาติ, ปทุมธานี, 153 น.

Caldwell, S., 2010, Statistics Unplugged, 3rd
Ed., Thomson Wadsworth, California.

Cohen, B.H., 2008, Explaining Psychological
Statistics, 3rd Ed., John Wiley and Sons,
Inc., New Jersey.

Devore, J.L., 2012, Probability and Statistics for
Engineering and the Sciences, 8th Ed.,
Wadsworth Cengage Learning, Inc.,
California.

Hogg, R.V., and Tanis, E.A., 2010, Probability
and Statistical Inference, 8th Ed., Upper
Saddle River, New Jersey.

Howell, D.C., 2012, Statistical Methods for
Psychology, 8th Ed., Wadsworth Cengage
Learning, Inc., California.

Pryczak, F., 1995, Making Sense of Statistics:
A Conceptual Overview, In Caldwell, S.,
2010, Statistics Unplugged, 3rd Ed.,
Wadsworth Cengage Learning, Inc.,
California.

Salkind, N.J., 2008, Statistics for People Who
(Think They) Hate Statistics, 3rd Ed.,
Sage Publications, Inc., California.