

บทที่ 2

ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

2.1 ทฤษฎีที่เกี่ยวข้อง

2.1.1 การวิเคราะห์ความแปรปรวน (Analysis of Variance)

จุดประสงค์ของการวิเคราะห์ความแปรปรวน

1. เพื่อทดสอบความแตกต่างของค่าเฉลี่ยประชากรมากกว่าสองชุด แต่ในการศึกษาการทดสอบสมมติเกี่ยวกับค่าเฉลี่ย ค่าสัดส่วน หรือค่าแปรปรวนของประชากร 2 ประชากร โดยการสุ่มตัวอย่าง 2 ชุดที่เป็นอิสระกันที่มีการแจกแจงปรกตินั้น การวิเคราะห์ดังกล่าวนั้นจะใช้สำหรับการเปรียบเทียบข้อมูลเพียง 2 ชุด เช่น ต้องการเปรียบเทียบยอดขายเฉลี่ยของพนักงานขาย 2 คน แต่ในความเป็นจริง บริษัทที่มีพนักงานขายมากกว่า 2 คน และต้องการเปรียบเทียบยอดขายเฉลี่ยของพนักงานเหล่านี้ จะเกิดปัญหาไม่สามารถเปรียบเทียบค่าเฉลี่ยหลายๆค่าพร้อมกันได้ หรือ อาจทำได้โดยต้องทำการทดสอบเปรียบเทียบทีละคู่ โดยใช้การทดสอบ Z หรือ T ซึ่งมีผลต่อความสามารถในการควบคุมความน่าจะเป็นที่จะเกิดความผิดพลาดประเภทที่ 1 ที่เกิดขึ้นในแต่ละการทดสอบเปรียบเทียบทีละคู่ ดังนั้น นักสถิติได้เสนอแนะวิธีที่จะใช้วิเคราะห์ปัญหาเช่นนี้ให้ได้ผลดี เรียกว่า การวิเคราะห์ความแปรปรวน ซึ่งสามารถทำการเปรียบเทียบความแตกต่างของค่าเฉลี่ยหลายๆ ค่าพร้อมกันได้

2. เพื่อแยกส่วน (partition) ความแปรปรวนทั้งหมดออกตามสาเหตุ โดยแยกเป็นความแปรปรวนที่เกิดจากความแตกต่างระหว่างกลุ่ม (Between-group Variation) และความแปรปรวนภายในกลุ่ม (Within-group Variation) ซึ่งเกิดจากหน่วยทดลองต่างกัน จึงมักเรียกความแปรปรวนส่วนนี้ว่า ความผิดพลาดจากการทดลอง (Experimental Error) หรือความผิดพลาดจากโอกาส (by chance)

หลักการวิเคราะห์ความแปรปรวน

การเปรียบเทียบค่าเฉลี่ยของประชากรมากกว่า 2 ประชากร จะไม่ทำการเปรียบเทียบค่าเฉลี่ยของประชากรทีละคู่ โดยใช้ตัวทดสอบ $Z - test$ หรือ $t - test$ เนื่องจากทำให้เกิดความผิดพลาดประเภทที่ 1 (type I error) เพิ่มมากขึ้น เนื่องจากการเปรียบเทียบค่าเฉลี่ยทีละคู่หลายๆ ครั้ง ดังนั้น ควรทำการเปรียบเทียบค่าเฉลี่ยของประชากรมากกว่า 2 กลุ่มเพียงครั้งเดียวโดยตัวสถิติที่ใช้คือการวิเคราะห์ความแปรปรวน (ANOVA F - test)

สถิติทดสอบเอฟเป็นตัวสถิติที่ใช้ในการทดสอบความแตกต่างของประชากร 2 กลุ่ม ได้แก่ (1) ความแปรปรวนระหว่างกลุ่ม (between group) และ (2) ความแปรปรวนภายในกลุ่ม (within group) โดยมีหลักการของการวิเคราะห์ความแปรปรวน คือ ถ้าความแปรปรวนระหว่างกลุ่มเท่ากับความแปรปรวนภายในกลุ่ม หมายความว่า ค่าเฉลี่ยของแต่ละประชากรที่มีมากกว่า 2 กลุ่มนั้นไม่มีความแตกต่างกันทางสถิติ ($\mu_1 = \mu_2 = \mu_3 = \dots = \mu_k$) แต่ถ้าความแปรปรวนระหว่างกลุ่มมีค่ามากกว่าความแปรปรวนภายในกลุ่ม หมายความว่า ค่าเฉลี่ยของแต่ละประชากรนั้นมีความแตกต่างกัน ยิ่งถ้าความแปรปรวนระหว่างกลุ่มมีค่ามากกว่าความแปรปรวนภายในกลุ่มเท่าใด (สถิติทดสอบเอฟมีค่าสูงมาก) ก็จะทำให้ค่าเฉลี่ยของแต่ละประชากรแตกต่างกันมากเท่านั้น (กวี สุจิตฺติ, 2546)

หลักการของการวิเคราะห์ความแปรปรวน คือ การแยกความผันแปรทั้งหมดออกเป็น ส่วนๆ ที่เกิดขึ้นในการทดลอง โดยที่ความผันแปรรวมของข้อมูลที่วัดโดยผลบวกกำลังสอง (sum square total : SST) สามารถแบ่งออกได้เป็นสองส่วน คือ ส่วนแรก เกิดจากผลบวกกำลังสองเนื่องจากความแตกต่างระหว่างค่าเฉลี่ยของทรีตเมนต์กับค่าเฉลี่ยหลัก (sum square treatment : $SSTr$) ซึ่งใช้วัดความแตกต่างระหว่างค่าเฉลี่ยของทรีตเมนต์แต่ละกลุ่ม จึงเรียกว่า “ผลบวกกำลังสองเนื่องจากทรีตเมนต์” และส่วนที่สอง เกิดจากผลบวกกำลังสองอันเนื่องมาจากความแตกต่างของค่าสังเกตภายในทรีตเมนต์หนึ่งๆ กับค่าเฉลี่ยของทรีตเมนต์นั้นๆ (sum square error : SSE) ซึ่งเป็นผลเนื่องมาจากความผิดพลาดสุ่ม จึงเรียกว่า “ผลบวกกำลังสองเนื่องจากความผิดพลาด” (กมล บุษบา, 2551)

ตัวแบบสำหรับการวิเคราะห์ความแปรปรวน คือ

$$y_{ij} = \mu_i + \varepsilon_{ij}; i = 1, \dots, k, j = 1, \dots, n_i$$

โดยที่ y_{ij} คือ ค่าสังเกตที่ j จากทรีตเมนต์ที่ i ($i = 1, 2, 3, \dots, k$ ทรีตเมนต์และ $j = 1, 2, 3, \dots, n_i$ ซ้ำ)

μ_i คือ ค่าเฉลี่ยประชากรของทรีตเมนต์กลุ่มที่ i

ε_{ij} คือ ความผิดพลาดสุ่ม

ผลบวกกำลังสอง คือ

$$\sum_{i=1}^k \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{..})^2 = \sum_{i=1}^k \sum_{j=1}^{n_i} (\bar{y}_{i.} - \bar{y}_{..})^2 + \sum_{i=1}^k \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{i.})^2$$

$$SST = SStr + SSE$$

สมมติฐานในการทดสอบ

$$H_0: \mu_1 = \mu_2 = \mu_3 = \dots = \mu_k$$

$$H_1: \text{มี } \mu_i \text{ อย่างน้อยหนึ่งค่าที่แตกต่างจากค่าอื่น (} i = 1, \dots, k \text{)}$$

ข้อตกลงเบื้องต้นเกี่ยวกับการวิเคราะห์ความแปรปรวน

ในการวิเคราะห์ความแปรปรวนมีข้อตกลงเบื้องต้นในทอมของความผิดพลาดสุ่ม ε_{ij} ดังนี้

1. ความผิดพลาดสุ่มเป็นอิสระกัน
2. ความผิดพลาดมีการแจกแจงปกติ
3. ความผิดพลาดมีความแปรปรวนเท่ากันและคงที่

กล่าวคือ ผิดพลาดสุ่ม ε_{ij} แต่ละตัวต่างก็เป็นตัวแปรสุ่มที่เป็นอิสระกัน โดยที่มีการแจกแจงปกติที่มีค่าเฉลี่ยเป็นศูนย์และความแปรปรวนเท่ากัน คือ $\varepsilon_{ij} \stackrel{iid}{\sim} N(0, \sigma^2)$

ตารางวิเคราะห์ความแปรปรวน

แหล่งความผันแปร	ระดับขั้นความเสรี	ผลบวกกำลังสอง	กำลังสองเฉลี่ย	เอฟ
ระหว่างทรีตเมนต์	$k - 1$	$SSTr$	$MSTr$	$F = \frac{MSTr}{MSE}$
ความผิดพลาด	$N - k$	SSE	MSE	
รวม	$N - 1$	SST		

การตรวจสอบสมมติฐานเกี่ยวกับข้อตกลงเบื้องต้นของการวิเคราะห์ความแปรปรวน

การตรวจสอบความเป็นอิสระกัน

ทำได้โดยการเลือกตัวอย่างแบบสุ่ม คือเลือกแบบไม่เจาะจง หรือโดยการกำหนดหน่วยทดลองให้แต่ละทรีตเมนต์อย่างสุ่มหรือไม่เจาะจง ซึ่งจากวิธีดังกล่าวจะทำให้มั่นใจได้ว่าตัวอย่างถูกเลือกมาอย่างอิสระกัน

การตรวจสอบการแจกแจงปกติ

พิจารณาจากกราฟ ฮิสโตแกรม หรือใช้หลักการทางสถิติทดสอบ

การตรวจสอบความเท่ากันของความแปรปรวน

การทดสอบความเท่ากันของความแปรปรวนทำได้โดย พิจารณาจากกราฟ ฮิสโตแกรม หรือใช้หลักการทางสถิติทดสอบ

เดตัน (Dayton, 1970) กล่าวว่า ในแบบแผนการวิจัยแบบสุ่มสมบูรณ์นั้น ความแปรปรวนในแต่ละกลุ่ม จะประมาณได้จากความแปรปรวนของประชากรโดยทำการตรวจสอบความเท่ากันของความแปรปรวน คือ $\sigma_1^2 = \sigma_2^2 = \sigma_3^2 = \dots = \sigma_k^2$ เมื่อ σ_i^2 แทน ความแปรปรวนของประชากรกลุ่มที่ i แต่เมื่อทำการทดสอบกับกลุ่มตัวอย่างจะได้รับความแปรปรวนกลุ่มตัวอย่าง คือ $\hat{\sigma}_1^2, \hat{\sigma}_2^2, \dots, \hat{\sigma}_i^2$ เป็นค่าประมาณความแปรปรวนประชากร สำหรับวิธีการตรวจสอบความเท่ากันของความแปรปรวนที่นิยมใช้กันอยู่อย่างแพร่หลายมีอยู่ 3 วิธี ได้แก่

1. วิธีของ บาร์ทเลท (Bartlett's test)
2. วิธีของ ฮาร์ดเลย์ (Hartlay's test)
3. วิธีของ เลวิน (Levene's test)

วิธีที่ใช้สำหรับการตรวจสอบเกี่ยวกับความแปรปรวนของประชากร k ประชากร ($k > 2$) มีหลายวิธีด้วยกัน แต่ในการศึกษาครั้งนี้ใช้วิธีของเลวิน (Levene's test)

เลวิน (Levene) ได้นำเสนอวิธีการทดสอบความเท่ากันของความแปรปรวนของประชากร ซึ่งมีวิธีการคำนวณคล้ายกับการวิเคราะห์ความแปรปรวนทางเดียว โดยการแทนค่าสังเกตแต่ละค่าด้วยค่าสัมบูรณ์ของส่วนเบี่ยงเบนระหว่างค่าสังเกตแต่ละค่ากับค่ากลางของกลุ่มตัวอย่างนั้นๆ ซึ่งเป็นตัวทดสอบที่มีความแรงต่อลักษณะการแจกแจงประชากรที่ไม่เป็นปกติ (นิภาพร ขำสอาด, 2552)

สมมติฐานในการทดสอบคือ

$$H_0: \sigma_1^2 = \sigma_2^2 = \dots = \sigma_k^2$$

$$H_1: \sigma_i^2 \text{ อย่างน้อย 1 ค่าที่แตกต่างกันจากค่าอื่น } i = 1, 2, \dots, k$$

คำนวณสถิติทดสอบได้จากสูตร

$$L = \frac{\sum_{i=1}^k n_i (Z_{i.} - Z_{..})^2 / (k-1)}{\sum_{i=1}^k \sum_{j=1}^{n_i} (Z_{ij} - Z_{i.})^2 / (N-k)} \quad ; \quad df = (k-1, N-k)$$

$$Z_{ij} = |y_{ij} - \bar{y}_{i.}| \quad , \quad Z_{i.} = \sum_{j=1}^{n_i} Z_{ij} / n_i \quad , \quad Z_{..} = \sum_{i=1}^k \sum_{j=1}^{n_i} Z_{ij} / N$$

เมื่อ	L	คือ	ค่าสถิติทดสอบ
	k	คือ	จำนวนทรีตเมนต์
	N	คือ	จำนวนค่าสังเกตทั้งหมด
	n_i	คือ	จำนวนค่าสังเกตในทรีตเมนต์กลุ่มที่ i
	y_{ij}	คือ	ค่าสังเกตที่ j ในทรีตเมนต์กลุ่มที่ i
	$\bar{y}_{i.}$	คือ	ค่ากลางของทรีตเมนต์กลุ่มที่ i

เกณฑ์การตัดสินใจเราจะปฏิเสธสมมติฐาน H_0 ถ้า $L > F_{\alpha, (k-1, N-k)}$ จากตารางการแจกแจงเอฟ ที่ระดับนัยสำคัญ α และองศาความอิสระ $k-1, (N-k)$

2.1.2 การเปรียบเทียบพหุคูณ (Multiple Comparison)

จากการวิเคราะห์ความแปรปรวน (Analysis of Variance ; ANOVA) ของข้อมูลในแผนการทดลองต่างๆ โดยใช้การทดสอบแบบ F-test มีจุดประสงค์เพียงเพื่อตรวจสอบว่าค่าเฉลี่ยของทรีตเมนต์ที่มีทั้งหมดในแผนการทดลองนั้นว่ามีความแตกต่างกันทางสถิติหรือไม่ แต่ไม่สามารถบ่งบอกให้ทราบว่าค่าเฉลี่ยของทรีตเมนต์ใดบ้างที่แตกต่างกัน เช่น การวิเคราะห์ความแปรปรวน สรุปว่า ปฏิเสธสมมติฐานว่าง ($H_0: \mu_1 = \mu_2 = \mu_3 = \dots = \mu_k$) หรือยอมรับสมมติฐานทางเลือก (H_a : มีค่าเฉลี่ยทรีตเมนต์ของประชากรอย่างน้อย 1 คู่ที่แตกต่างกัน) แสดงว่าค่าเฉลี่ยของทรีตเมนต์นั้นมีความแตกต่างกันทางสถิติ ทั้งนี้เพราะว่าการทดสอบแบบเอฟนี้เป็นการทดสอบไปพร้อมๆ กันว่า โดยเฉลี่ยแล้วทุกๆ ทรีตเมนต์แตกต่างกันหรือไม่ ดังนั้นจึงต้องทำการตรวจสอบความแตกต่างระหว่างทรีตเมนต์ ที่เรียกว่าการเปรียบเทียบพหุคูณ (Multiple Comparison) เพื่อให้การสรุปผลชัดเจนว่าค่าเฉลี่ยของทรีตเมนต์คู่ใดบ้างที่แตกต่างกัน เพื่อสะดวกในการนำไปใช้ประโยชน์ต่อไป

การตรวจสอบความแตกต่างของค่าเฉลี่ยทรีตเมนต์มีหลายวิธีขึ้นอยู่กับวัตถุประสงค์ของการเปรียบเทียบ ซึ่งแต่ละวิธีการเปรียบเทียบพหุคูณนั้นจะมีความเหมาะสมกับแต่ละลักษณะของงานและลักษณะของทรีตเมนต์ โดยสามารถแบ่งวิธีการเปรียบเทียบออกเป็น 2 ประเภท คือ

1. การเปรียบเทียบค่าเฉลี่ยของทรีตเมนต์ที่ไม่มีการวางแผนไว้ล่วงหน้า

วิธีการตรวจสอบประเภทนี้ผู้ทดลองไม่ได้คาดการณ์ไว้ก่อนว่าค่าเฉลี่ยของทรีตเมนต์ใดบ้างที่มีความแตกต่างกัน ทั้งนี้อาจเนื่องมาจากไม่มีรายละเอียด หรือมีข้อมูลที่เกี่ยวข้องกับทรีตเมนต์น้อย ไม่เพียงพอต่อการคาดการณ์ หรืออาจเนื่องมาจากทรีตเมนต์นั้นไม่สามารถจัดแบ่งเป็นกลุ่มหรือเป็นพวกได้อย่างชัดเจน จึงต้องทำการวิเคราะห์ความแปรปรวน (ANOVA) ก่อนเมื่อปฏิเสธสมมติฐานว่าง (H_0) แล้วจึงทำการเปรียบเทียบพหุคูณต่อไป ซึ่งสามารถเปรียบเทียบค่าเฉลี่ยทรีตเมนต์ทุกคู่ที่เป็นไปได้ทั้งหมด ดังนั้นการเปรียบเทียบประเภทนี้จึงมีรูปแบบที่แน่นอน (conventional comparison) ซึ่งมีหลายวิธี เช่น

- 1) Least significant difference (LSD)
- 2) Tukey's Honestly significant difference (HSD)
- 3) Duncan's new multiple range test (DMRT)
- 4) Student - Newman - Keuls procedure (S-N-K)

นอกจากนี้ถ้าต้องการเปรียบเทียบเพียงบางคู่ หรือต้องการเปรียบเทียบทุกคู่ของการเปรียบเทียบที่เป็นไปได้ทั้งหมด (For comparison any all possible contrasts between treatment means) สามารถใช้วิธีการเปรียบเทียบค่าเฉลี่ยทรีตเมนต์ของ Scheffe's method

2. การเปรียบเทียบค่าเฉลี่ยของทรีตเมนต์ที่มีการวางแผนไว้ล่วงหน้า

เป็นวิธีการเปรียบเทียบพหุคูณที่ไม่ต้องมีการวิเคราะห์ความแปรปรวน (ANOVA) เนื่องจากผู้ทดลองได้วางแผนไว้ล่วงหน้าแล้วว่าจะทำการทดสอบค่าเฉลี่ยของทรีตเมนต์คู่ใดบ้าง ทั้งนี้เนื่องจากผู้ทดลองพอที่จะทราบรายละเอียดเกี่ยวกับทรีตเมนต์ดีพอ อาทิ ทราบถึงลักษณะของทรีตเมนต์ว่าเป็นลักษณะปริมาณหรือคุณภาพ, โครงสร้างของทรีตเมนต์ว่าสามารถแบ่งเป็นกลุ่มได้หรือไม่ และทรีตเมนต์มาตรฐานหรือตัวเปรียบเทียบ คือ การทดลองที่มีการเปรียบเทียบทรีตเมนต์ใหม่กับทรีตเมนต์ควบคุมเป็นต้น

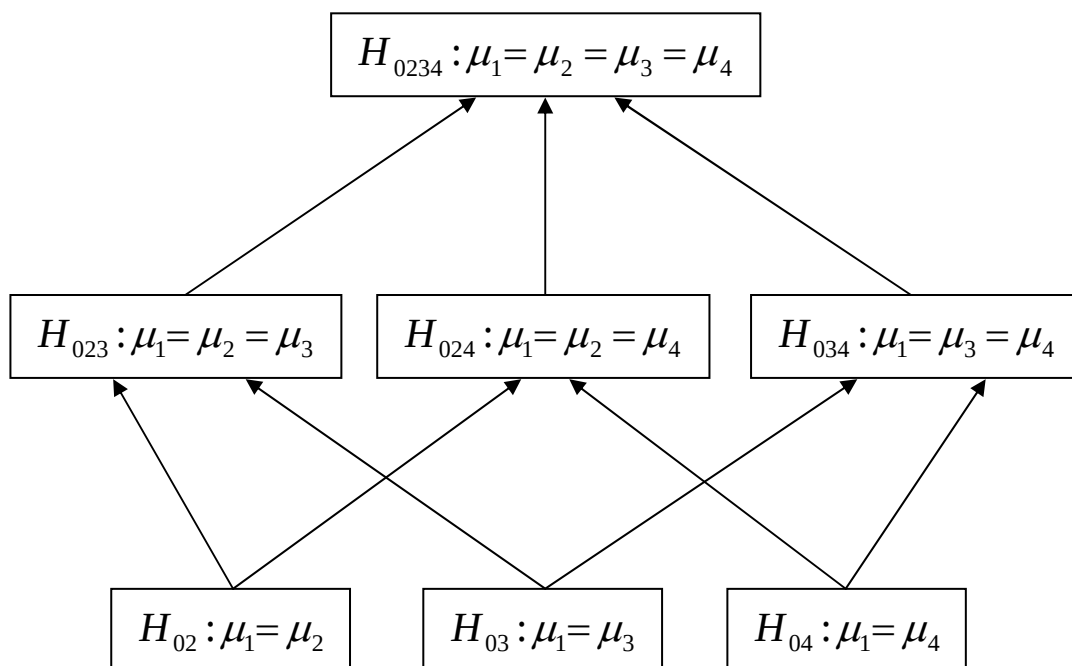
ดังนั้นการเปรียบเทียบประเภทนี้ พบว่า จำนวนคู่ของการเปรียบเทียบค่าเฉลี่ยทรีตเมนต์ต่างๆ จะมีจำนวนน้อยกว่าการเปรียบเทียบที่เป็นไปได้ทั้งหมด ซึ่งวิธีการเปรียบเทียบสามารถแบ่งได้ดังนี้

- 1) Orthogonal polynomial เป็นการเปรียบเทียบพหุคูณว่ามีแนวโน้มอย่างไร เมื่อทรีตเมนต์มีลักษณะเป็นปริมาณ
- 2) Contrasts หรือ comparison เป็นการเปรียบเทียบพหุคูณ เมื่อทรีตเมนต์มีลักษณะเป็นคุณภาพ ซึ่งเป็นการเปรียบเทียบระหว่างกลุ่ม 2 กลุ่ม
- 3) Least significant difference (LSD) และ Dunnett's method เมื่อต้องการเปรียบเทียบทรีตเมนต์เพียงบางคู่ โดยในการทดลองนั้นมีทรีตเมนต์ควบคุม (กวี สุจิตฺติ, 2546)

2.1.3 วิธีการทดสอบแบบปิด (Closed test)

ถ้าเราต้องการทดสอบสมมติฐาน $H_{01}, H_{02}, \dots, H_{0k}$ เมื่อ H_{0k} คือ สมมติฐานว่างของการเปรียบเทียบค่าเฉลี่ยทรีตเมนต์ที่ k เช่น ต้องการเปรียบเทียบค่าเฉลี่ยของทรีตเมนต์กับทรีตเมนต์ควบคุมที่มี 3 ทรีตเมนต์ และทรีตเมนต์ควบคุม 1 ทรีตเมนต์ จะมีสมมติฐานว่าง $H_{0234}:\mu_1=\mu_2=\mu_3=\mu_4$ ดังนั้นจะมีสมมติฐานเชิงเดียวที่ทดสอบความแตกต่างระหว่างค่าเฉลี่ยทรีตเมนต์กับทรีตเมนต์ควบคุม จำนวน 3 สมมติฐานการทดสอบ คือ $H_{02}:\mu_1=\mu_2$, $H_{03}:\mu_1=\mu_3$ และ $H_{04}:\mu_1=\mu_4$ เมื่อกำหนดให้ทรีตเมนต์ที่ 1 เป็นทรีตเมนต์ควบคุม วิธีการทดสอบแบบปิด มีขั้นตอนดังนี้

1. ในการทดสอบแต่ละสมมติฐาน H_{02}, H_{03} และ H_{04} จะใช้ระดับนัยสำคัญของการทดสอบ (α) ที่เหมาะสม
2. สร้างเซตของสมมติฐานอินเตอร์เซกชันที่เป็นไปได้ทั้งหมดระหว่าง H_{02}, H_{03} และ H_{04} ซึ่งในที่นี้จะได้สมมติฐานคือ $H_{023}, H_{024}, H_{034}$ และ H_{0234}



3. การทดสอบแต่ละสมมติฐานอินเตอร์เซกชันจะต้องใช้ระดับนัยสำคัญที่เหมาะสม ซึ่งในการทดสอบอาจใช้ F-test หรือวิธีอื่นๆ ที่ใช้สำหรับการทดสอบสมมติฐานอินเตอร์เซกชัน ซึ่งแต่ละวิธีของการทดสอบแบบปิดจะให้ผลลัพธ์ที่แตกต่างกัน
4. ในการทดสอบจะปฏิเสธสมมติฐาน H_{0k} โดยการควบคุมอัตราความผิดพลาดสูงสุดต่อการทดลอง (Maximum Experimentwise Error Rate: MEER) ที่เป็นไปตามเงื่อนไข 2 ข้อ คือ
 - 1) การทดสอบสมมติฐานเชิงเดียว H_{0k} มีนัยสำคัญ
 - 2) การทดสอบทุกๆ สมมติฐานอินเตอร์เซกชันต้องประกอบด้วยสมมติฐานเชิงเดียว H_{0k} ที่มีนัยสำคัญ

กรณีที่เราสงสัยทดสอบความแตกต่างระหว่างค่าเฉลี่ยทรีตเมนต์กับทรีตเมนต์ควบคุม $H_{0234}:\mu_1=\mu_2=\mu_3=\mu_4$ ถ้าผลการทดสอบสมมติฐานเชิงเดียวนั้นมีนัยสำคัญ ก็จะส่งผลให้การทดสอบสมมติฐาน ดังกล่าวมีนัยสำคัญด้วย

2.1.4 การสุ่มซ้ำโดยวิธีบูทสแตรป (Bootstrap Resampling)

2.1.4.1 การสุ่มซ้ำแบบคินที่โดยวิธีบูทสแตรป (Independent bootstrap resampling)

พิจารณาตัวแปรสุ่มที่มีรูปแบบการแจกแจงที่เหมือนกันและเป็นอิสระต่อกัน วิธีการสุ่มแบบคินที่โดยวิธีบูทสแตรป มีดังนี้

กำหนดให้ $\{X_{n,j}^*, 1 \leq j \leq m\}$ เป็นตัวอย่างสุ่มจำนวน m ค่าที่ได้จากการสุ่มแบบคินที่จากค่าสังเกตหรือข้อมูลตัวอย่าง $X_1, X_2, \dots, X_n$ จำนวน n ค่า ดังนั้น ตัวอย่างสุ่มแต่ละ X_i , $i = 1, \dots, n$ จะถูกเลือกด้วยความน่าจะเป็น $\frac{1}{n}$ เป็นจำนวน m ค่า

เมื่อเรามีข้อมูลตัวอย่างสุ่มขนาด n จากประชากร ซึ่งจากตัวอย่างขนาด n นี้ เราจะทำการสุ่มแบบคินที่ซึ่งข้อมูลที่ถูกเลือก จะถูกใส่คืนไปแล้วเลือกออกมาใหม่จนครบจำนวน m ค่า เราเรียกตัวอย่างสุ่มจำนวน m ที่ได้นี้ว่า ตัวอย่างสุ่มแบบบูทสแตรปที่เป็นอิสระกัน (Independent bootstrap sample) เช่น เรามีข้อมูลตัวอย่างสุ่มจากประชากรขนาด $n = 4$ คือ $\{2, 4, 5, 7\}$ เราจะสุ่มตัวอย่างขนาด $m = 6$ จากตัวอย่างสุ่มข้างต้น โดยตัวอย่างสุ่มแต่ละตัวที่ถูกเลือกจะถูกใส่กลับคืนแล้วทำการสุ่มเลือกออกมาใหม่จนครบ 6 ค่า เช่น ชุดข้อมูลตัวอย่างที่เราได้ คือ $\{7, 7, 5,$

2, 2, 7} ซึ่งเรียกดัวย่างที่ได้นี้ว่าตัวอย่างสุ่มแบบบูทสเตรปที่เป็นอิสระกัน (Independent bootstrap sample) (Tosasukul, J., 2007)

2.1.4.2 การสุ่มซ้ำแบบไม่คืนที่โดยวิธีบูทสเตรป (Dependent bootstrap resampling)

การสุ่มแบบไม่คืนที่โดยวิธีบูทสเตรปนี้เป็นวิธีที่ลดความผิดพลาดที่เกิดจากการประมาณค่าและให้ช่วงความเชื่อมั่นที่ดี เรายินยอมให้ตัวแปรสุ่มที่ได้จากการสุ่มแบบไม่คืนที่โดยวิธีบูทสเตรป มีดังนี้

กำหนดให้ $\{m, m \geq 1\}$ และ $\{c, c \geq 1\}$ เป็นจำนวนเต็มบวก โดยที่ $m \leq nc$ สำหรับทุกๆ $n \geq 1$ ดังนั้น เราจะกำหนดให้ $\{X'_{cn,j}, 1 \leq j \leq m\}$ เป็นตัวอย่างสุ่มจำนวน m ค่าที่ได้จากการสุ่มแบบไม่คืนที่จากค่าสังเกตหรือข้อมูลตัวอย่าง $X_1, X_2, \dots, X_n$ ที่ได้ทำการคัดลอกจำนวน c ครั้ง

เมื่อเรามีข้อมูลตัวอย่างสุ่มขนาด n จากประชากร ซึ่งจากตัวอย่างขนาด n นี้ เราจะทำการคัดลอกข้อมูลตัวอย่างจำนวน c ชุด ดังนั้นจะมีจำนวนข้อมูล nc จำนวน หลังจากนั้นทำการสุ่มแบบไม่คืนที่จำนวน m ค่า ซึ่งข้อมูลที่ถูกเลือกออกมาแล้วจะไม่ถูกเลือกออกมาอีก โดยที่เซตของตัวอย่างสุ่มที่ได้นี้ จะมีจำนวนน้อยกว่าหรือเท่ากับ nc , ($m \leq nc$) เราเรียกชุดตัวอย่างสุ่มจำนวน m ค่าที่ได้นี้ว่า ตัวอย่างสุ่มแบบบูทสเตรปที่ไม่เป็นอิสระกัน (Dependent bootstrap sample) เช่น เรามีข้อมูลตัวอย่างสุ่มจากประชากรขนาด $n=4$ คือ $\{2, 4, 5, 7\}$ สมมติเราต้องการจะสุ่มตัวอย่างขนาด $m=6$ จากการคัดลอกตัวอย่างสุ่มจากการประชากรข้างต้นจำนวน 3 ชุด ดังนั้นข้อมูลที่ได้จากการคัดลอกคือ $\{2, 4, 5, 7, 2, 4, 5, 7, 2, 4, 5, 7\}$ เมื่อได้ข้อมูลแล้วทำการสุ่มแบบไม่คืนที่จนครบ 6 ค่า ($m \leq nc = (4)(3) = 12$) เช่น ชุดข้อมูลตัวอย่างที่เราได้ คือ $\{4, 7, 4, 5, 5, 2\}$ ซึ่งเรียกดัวย่างที่ได้นี้ว่าตัวอย่างสุ่มแบบบูทสเตรปที่ไม่เป็นอิสระกัน (Dependent bootstrap sample) (Tosasukul, J., 2007)

2.1.5 การกำหนดระดับความแตกต่างของความแปรปรวน

กำหนดความแปรปรวนของประชากรแต่ละกลุ่มเป็น 3 ระดับคือ แตกต่างกันน้อย แตกต่างกันปานกลางและแตกต่างกันมาก กำหนดตามวิธีของเกมส์ พรอบเบอทและวังก์เรอ (Games, P.A., Probert, D.A. and Winkler, H.B., 1972) โดยใช้ค่าอนเซนทรลิตี้ พารามิเตอร์ ϕ (noncentrality parameter) เป็นเกณฑ์วัดความแตกต่างความแปรปรวนของประชากร ดังต่อไปนี้

1. อัตราส่วนของความแปรปรวนแตกต่างกันน้อย ถ้า $0 \leq \phi < 1.5$
2. อัตราส่วนของความแปรปรวนแตกต่างกันปานกลาง ถ้า $1.5 \leq \phi < 3.0$
3. อัตราส่วนของความแปรปรวนแตกต่างกันมาก ถ้า $\phi \geq 3.0$

สูตรที่ใช้ในการคำนวณค่า ϕ คือ

$$\phi^2 = \sum_{i=1}^k \frac{(\sigma_i^2 - \sigma^2)^2}{k\sigma_1^2}$$

- เมื่อ σ_1^2 คือ ความแปรปรวนของประชากรที่มีค่าน้อยที่สุด
 σ_i^2 คือ ความแปรปรวนของประชากรกลุ่มที่ i ($i=1, 2, \dots, k$)
 σ^2 คือ ค่าเฉลี่ยเรขาคณิตของความแปรปรวนของประชากร k กลุ่ม

$$\text{โดยที่ } \sigma^2 = \left(\prod_{i=1}^k \sigma_i^2 \right)^{\frac{1}{k}}$$

2.1.6 การทดสอบความสามารถในการควบคุมความน่าจะเป็นที่จะเกิดความผิดพลาดประเภทที่ 1

ในการทดสอบว่าสถิติทดสอบแต่ละตัวสามารถควบคุมความน่าจะเป็นที่จะเกิดความผิดพลาดประเภทที่ 1 นั้น จะใช้การทดสอบทวินาม โดยกำหนดระดับนัยสำคัญในการทดสอบเท่ากับ 0.05 โดยมีวิธีการ ดังนี้

สมมติฐานการทดสอบ คือ

$$H_0 : \alpha = \alpha_0$$

$$H_1 : \alpha \neq \alpha_0$$

สถิติทดสอบ คือ

$$Z = \frac{\alpha_* - \alpha_0}{\sqrt{\frac{\alpha_0(1-\alpha_0)}{M}}}$$

กำหนดให้ α	แทน	ค่าความน่าจะเป็นของความผิดพลาดประเภทที่ 1
α_0	แทน	ระดับนัยสำคัญที่กำหนดในการศึกษาเท่ากับ 0.05
α_*	แทน	ค่าความน่าจะเป็นของความผิดพลาดประเภทที่ 1 ที่ได้จากการทดลอง ซึ่งหาได้จาก อัตราส่วนของจำนวน ครั้งของการปฏิเสธสมมติฐานว่างเมื่อสมมติฐานว่าง เป็นจริง เทียบกับจำนวนครั้งที่ทำการทดลอง
M	แทน	จำนวนครั้งที่ทำการทดลอง เท่ากับ 1,000 ครั้ง

เกณฑ์การพิจารณาความสามารถในการควบคุมความผิดพลาดประเภทที่ 1 เมื่อกำหนดระดับนัยสำคัญของการทดสอบทวินาม เท่ากับ 0.05 และจำนวนครั้งที่ทำการทดลอง (M) เท่ากับ 1,000 ครั้ง ซึ่งสถิติทดสอบที่ศึกษาครั้งนี้จะสามารถควบคุมความผิดพลาดประเภทที่ 1 ได้ ถ้า $Z \leq 1.645$ ซึ่งสามารถสรุปได้ดังนี้

$$Z \leq 1.645$$

$$Z = \frac{\alpha_* - 0.05}{\sqrt{\frac{0.05(1-0.05)}{1,000}}} \leq 1.645$$

$$\alpha_* \leq 0.061$$

เมื่อกำหนดระดับนัยสำคัญของการทดสอบ (α_0) เท่ากับ 0.05 สถิติทดสอบสามารถควบคุมความผิดพลาดประเภทที่ 1 ได้ ถ้า α_* มีค่าไม่เกิน 0.061

2.2 สถิติที่ใช้ในการศึกษา

2.2.1 สถิติทดสอบของดันเนตต์ (Dunnett's two sided test statistic)

การเปรียบเทียบพหุคูณโดยสถิติทดสอบของดันเนตต์

แผนการทดลองทั่วไปแบ่งทรีตเมนต์ออกเป็น 2 ประเภท คือ (1) ทรีตเมนต์ควบคุม (specified control treatment) และ (2) ทรีตเมนต์อื่นๆ ในการเปรียบเทียบพหุคูณโดยสถิติทดสอบของดันเนตต์นั้น มีวัตถุประสงค์เพื่อต้องการเปรียบเทียบค่าเฉลี่ยของทรีตเมนต์ต่างๆ กับทรีตเมนต์ควบคุม โดยมีวิธีการดังนี้

สมมติฐานการทดสอบ คือ

$$H_{0i} : \mu_c = \mu_i$$

$$H_{1i} : \mu_c \neq \mu_i$$

สถิติทดสอบ คือ

$$D_i = \frac{|\bar{y}_c - \bar{y}_i|}{\sqrt{MSE(1/n_c + 1/n_i)}}$$

โดยที่

$$MSE = \frac{\sum_{i=1}^k \sum_{j=1}^n (y_{ij} - \bar{y}_{i.})^2}{k(n-1)}$$

กำหนดให้ y_{ij} แทน ค่าสังเกตที่ j จากทรีตเมนต์ที่ i ($i = 1, 2, 3, \dots, k$ ทรีตเมนต์และ $j = 1, 2, 3, \dots, n$ ซ้ำ)

$\bar{y}_c$ แทน ค่าเฉลี่ยของทรีตเมนต์ควบคุม

$\bar{y}_i$ แทน ค่าเฉลี่ยของทรีตเมนต์กลุ่มที่ i โดยที่ $i = 1, 2, \dots, k-1$

n_c แทน ขนาดตัวอย่าง(จำนวนซ้ำ)ของทรีตเมนต์ควบคุม

n_i แทน ขนาดตัวอย่าง(จำนวนซ้ำ)ของทรีตเมนต์กลุ่มที่ i

MSE แทน ความคลาดเคลื่อนจากการวิเคราะห์ความแปรปรวน

ในกรณีที่ขนาดตัวอย่าง(จำนวนซ้ำ)ในแต่ละทรีตเมนต์มีจำนวนเท่ากัน ($n_c = n_i = n$) สามารถคำนวณค่า D_i ได้จากสูตร

$$D_i = \frac{|\bar{y}_c - \bar{y}_i|}{\sqrt{MSE(2/n)}}$$

เกณฑ์การตัดสินใจเราจะปฏิเสธสมมติฐานว่าง $H_{0i} : \mu_c = \mu_i$ เมื่อ $D_i > D_{\alpha, (k-1, k(n-1))}$ ที่ระดับนัยสำคัญ α โดยที่ $D_{\alpha, (k-1, k(n-1))}$ คือค่าวิกฤตที่ได้จากตารางดันเนตต์ (Dunnett's two-sided range distribution) ที่ระดับนัยสำคัญ α และองศาความอิสระ $k-1, k(n-1)$

2.2.2 วิธีสแต็ปดาวน์อินดิเพนเดนท บูทสเตรป มิน พี (Step-down Independent Bootstrap min P)

สูตรที่ใช้ในการหาค่าสถิติทดสอบสำหรับวิธีสแต็ปดาวน์อินดิเพนเดนท บูทสเตรป มิน พี คือ

$$\tilde{p}_{(i)} = \max_{m=1, \dots, i} \left\{ \Pr \left(\min_{l \in \{m, \dots, k\}} P_l \leq p_{(m)} \mid H_0 \right) \right\}$$

ในการศึกษาครั้งนี้ กำหนดให้ทรีตเมนต์ที่ทำการศึกษามี 3 ทรีตเมนต์ และมีทรีตเมนต์ควบคุม 1 ทรีตเมนต์ เมื่อกำหนดให้ทรีตเมนต์ที่ 1 เป็นทรีตเมนต์ควบคุม ($\mu_c = \mu_1$) สมมติฐานการทดสอบทั้งหมดในการเปรียบเทียบค่าเฉลี่ยรายคู่ระหว่างทรีตเมนต์ควบคุมกับทรีตเมนต์อื่นๆ มีดังนี้

$$\begin{aligned} H_{02} : \mu_1 &= \mu_2 & \text{vs} & & H_{12} : \mu_1 &\neq \mu_2 \\ H_{03} : \mu_1 &= \mu_3 & \text{vs} & & H_{13} : \mu_1 &\neq \mu_3 \\ H_{04} : \mu_1 &= \mu_4 & \text{vs} & & H_{14} : \mu_1 &\neq \mu_4 \end{aligned}$$

สถิติทดสอบ
$$t_i = \frac{\bar{y}_1 - \bar{y}_i}{\sqrt{MSE(2/n)}}, \text{ โดยที่ } i = 2, 3, 4$$

ปรับค่า p-value จากสูตร
$$\tilde{p}_{(i)} = \max_{m=2, \dots, i} \left\{ \Pr \left(\min_{l \in \{m, \dots, 4\}} P_l \leq p_{(m)} \mid H_0 \right) \right\}$$

โดยพิจารณาดังนี้

$$\tilde{p}_{(2)} = \max \left\{ \Pr \left(\min_{l \in \{2,3,4\}} P_l \leq p_{(2)} \right) \right\} = \Pr(\min(P_2, P_3, P_4) \leq p_{(2)})$$

$$\tilde{p}_{(3)} = \max \left\{ \Pr \left(\min_{l \in \{2,3,4\}} P_l \leq p_{(2)} \right), \Pr \left(\min_{l \in \{3,4\}} P_l \leq p_{(3)} \right) \right\}$$

$$\tilde{p}_{(4)} = \max \left\{ \Pr \left(\min_{l \in \{2,3,4\}} P_l \leq p_{(2)} \right), \Pr \left(\min_{l \in \{3,4\}} P_l \leq p_{(3)} \right), \Pr \left(\min_{l \in \{4\}} P_l \leq p_{(4)} \right) \right\}$$

โดยที่ $p_{(m)}$ คือ ค่า p-value ที่คำนวณได้จากการทดสอบสมมติฐานข้างต้นจำนวน 3 สมมติฐาน โดยคำนวณจากข้อมูลเริ่มต้นที่ได้ทำการจำลอง โดยเรียงค่าจากน้อยไปมาก เพื่อใช้ในการพิจารณาปรับค่า p-value

P_l คือ ค่า p-value ที่คำนวณได้จากการทดสอบสมมติฐานข้างต้นจำนวน 3 สมมติฐาน หลังจากการสุ่มซ้ำแบบคืนที่โดยวิธีบูทสเตรป

2.2.3 วิธีสแต็ปดาวน์ดิเพนเดนท บูทสเตรป มิน พี (Step-down Dependent Bootstrap min P)

สูตรที่ใช้ในการหาค่าสถิติทดสอบสำหรับวิธีการสแต็ปดาวน์ดิเพนเดนท บูทสเตรป มิน พี นั้นพิจารณาเหมือนกับวิธีการสแต็ปดาวน์ดิเพนเดนท บูทสเตรป มิน พี แต่จะแตกต่างกัน ที่การพิจารณาหา P_l ซึ่งจะคำนวณค่า p-value ที่ได้จากการทดสอบสมมติฐานข้างต้นจำนวน 3 สมมติฐาน หลังจากการสุ่มซ้ำแบบไม่คืนที่โดยวิธีบูทสเตรป ที่ได้จากการคัดลอกข้อมูลที่ได้จากการจำลองขั้นแรก ตามจำนวนการคัดลอกที่กำหนด

2.3 การแจกแจงที่ใช้ในการศึกษา

2.3.1 การแจกแจงปกติ (Normal Distribution)

กำหนดให้ X เป็นตัวแปรสุ่มต่อเนื่องที่มีการแจกแจงปกติ โดยมีพารามิเตอร์ μ และ σ^2 ดังนั้น ฟังก์ชันการแจกแจง (Distribution function) หรือ ฟังก์ชันความหนาแน่นความน่าจะเป็น (Probability density function) ของ X คือ

$$f(x; \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}, \quad -\infty \leq x < \infty, -\infty < \mu < \infty, \sigma^2 > 0$$

- ค่าเฉลี่ยของตัวแปรสุ่ม x
 $E(X) = \mu$
- ความแปรปรวนของตัวแปรสุ่ม x
 $\text{Var}(X) = \sigma^2$
- สัมประสิทธิ์ความเบ้
 $\gamma_1 = 0$
- สัมประสิทธิ์ความโด่ง
 $\gamma_2 = 3$

คุณสมบัติของการแจกแจงปกติ มีดังนี้

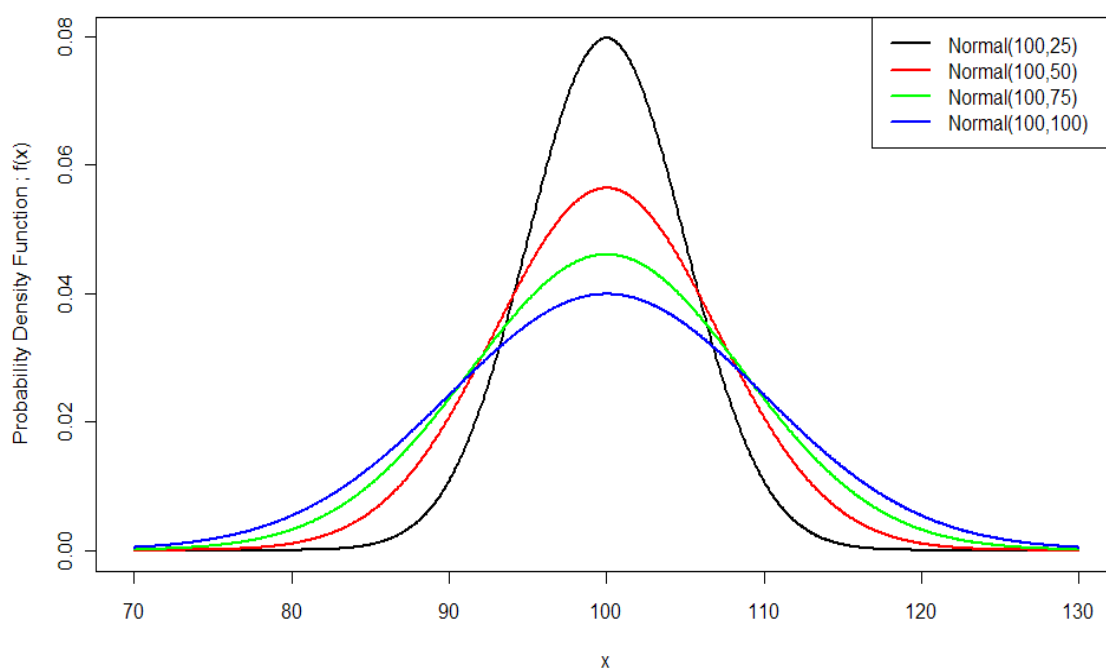
1. ฟังก์ชันความหนาแน่นความน่าจะเป็นมีลักษณะเป็นเส้นโค้งไม่มีความเบ้ หรือมีค่าความเบ้เท่ากับศูนย์
2. เส้นโค้งปกติเป็นรูปประฆังคว่ำสมมาตรรอบ $x = \mu$
3. ค่าเฉลี่ย มัธยฐาน และฐานนิยมมีค่าเท่ากันและอยู่ในตำแหน่งกลางเส้นโค้งในแนวตั้งฉาก
4. มีแกน X เป็นเส้นกำกับ (Asymptote) กล่าวคือ ปลายของเส้นโค้งทั้งสองข้างจะไม่ตัดแกน X
5. พื้นที่ใต้โค้งทั้งหมดคือความน่าจะเป็นของแซมเปิลสเปซ (Sample space) ซึ่งมีค่าเท่ากับ 1 หรือ 100%

6. พื้นที่ใต้โค้งระหว่างจุดที่เป็น $\pm 1\sigma$ มีประมาณ 68.26%, พื้นที่ใต้โค้งระหว่างจุดที่เป็น $\pm 2\sigma$ มีประมาณ 95.46% และพื้นที่ใต้โค้งระหว่างจุดที่เป็น $\pm 3\sigma$ มีประมาณ 99.74%

7. เมื่อการแจกแจงปกติมีค่า $\mu = 0$ และ $\sigma^2 = 1$ จะมีการแจกแจงที่เรียกว่าการแจกแจงปกติมาตรฐาน (Standard normal distribution)

ภาพที่ 2.1

แสดงฟังก์ชันความหนาแน่นของการแจกแจงปกติ



2.3.2 การแจกแจงล็อกนอร์มัล (Lognormal Distribution)

กำหนดให้ X เป็นตัวแปรสุ่มต่อเนื่องที่มีการแจกแจงล็อกนอร์มัล โดยมีพารามิเตอร์ μ และ σ^2 ดังนั้น ฟังก์ชันการแจกแจง (Distribution function) หรือ ฟังก์ชันความหนาแน่นความน่าจะเป็น (Probability density function) ของ X คือ

$$f(x; \mu, \sigma^2) = \frac{1}{x\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2}\left(\frac{\ln x - \mu}{\sigma}\right)^2}, \quad 0 \leq x < \infty, -\infty < \mu < \infty, \sigma^2 > 0$$

ลักษณะทั่วไปของตัวแปรสุ่ม X ที่มีการแจกแจงล็อกนอร์มัล

- ค่าเฉลี่ยของตัวแปรสุ่ม x

$$E(X) = \exp\left(\mu + \frac{\sigma^2}{2}\right)$$

- ความแปรปรวนของตัวแปรสุ่ม x

$$\text{Var}(X) = m^2 (\omega - 1)$$

- สัมประสิทธิ์ความเบ้

$$\gamma_1 = (\omega + 2) \omega \sqrt{\omega - 1}$$

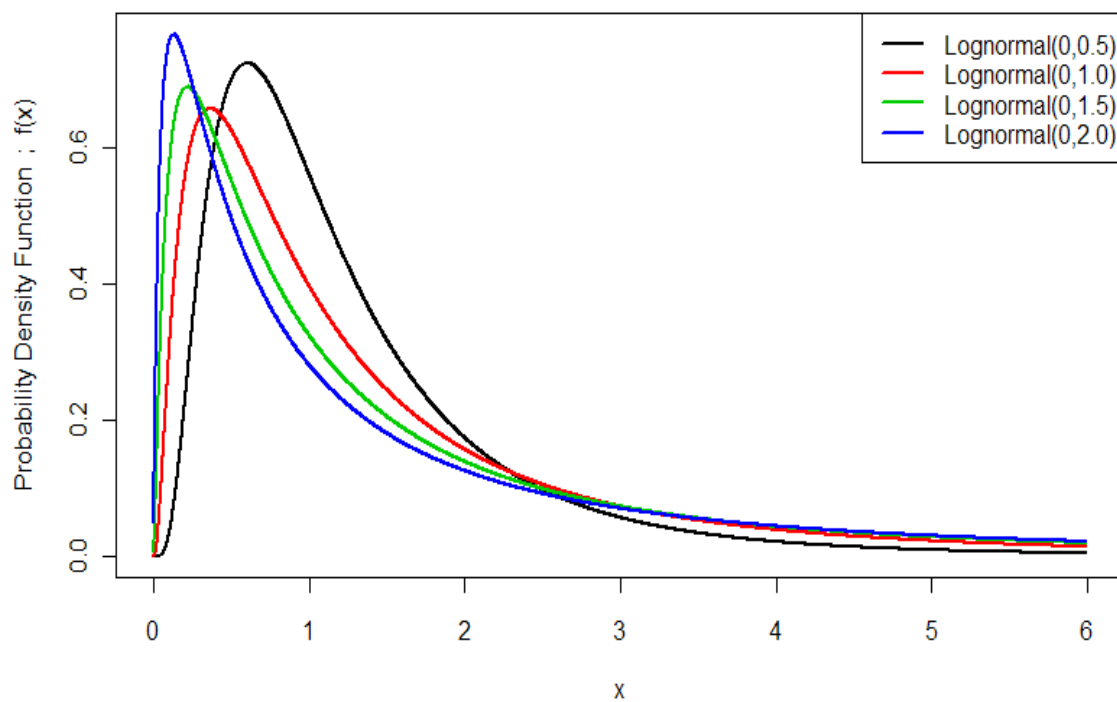
- สัมประสิทธิ์ความโด่ง

$$\gamma_2 = \omega^4 + 2\omega^3 + 3\omega^2 - 3$$

โดยที่ $m = \exp(\mu)$, $\omega = \exp(\sigma^2)$

ภาพที่ 2.2

แสดงฟังก์ชันความหนาแน่นของการแจกแจงล็อกนอร์มัล



2.4 เอกสารและงานวิจัยที่เกี่ยวข้อง

ในปี 2000 เวสต์ฟอลและโวลฟิงเกอร์ (Westfall and Wolfinger) ได้กล่าวถึง วิธีการทดสอบแบบปิด ซึ่งเป็นผลลัพธ์ในวิธีขั้นบันได (step-wise) ซึ่งเป็นวิธีที่ให้กำลังการทดสอบมากกว่า วิธีขั้นเดียว (single-step) ในสถานการณ์ที่คล้ายๆ กัน วิธีที่ง่ายที่จะกล่าวถึงคือ วิธีบอนเฟอร์โรนี (Bonferroni) ซึ่งเป็นตัวอย่างของวิธีขั้นเดียว ในแต่ละการเปรียบเทียบที่ต้องมีระดับนัยสำคัญ α/k เมื่อ k คือ จำนวนทรีตเมนต์ที่นำมาเปรียบเทียบ และ α คือ ความผิดพลาดประเภทที่ 1 ที่มากที่สุดของการเปรียบเทียบ k ทรีตเมนต์ ซึ่งวิธีขั้นบันไดจะใช้ค่าวิกฤตของระดับนัยสำคัญที่มากกว่า α/k นั้นหมายความว่า ที่ระดับนัยสำคัญต่างกันมาก วิธีขั้นบันไดจะมีกำลังการทดสอบมากกว่า

วิธีการขั้นตอนเดียว เป็นการเปรียบเทียบกันของค่าวิกฤตหลายๆ ค่า ที่ปรับรูปแบบทุกๆ สมมติฐาน โดยไม่คำนึงถึงลำดับค่าของสถิติทดสอบหรือไม่ปรับค่าพี (unadjusted p-value) ดังนั้น ในแต่ละสมมติฐานจะมาจากการใช้ค่าวิกฤต ซึ่งผลการทดสอบในแต่ละสมมติฐานจะเป็นอิสระจากสมมติฐานอื่นๆ ส่วนการปรับปรุงกำลังการทดสอบโดยการควบคุมความผิดพลาดประเภทที่ 1 อาจจะใช้วิธีการขั้นบันได ซึ่งสำหรับวิธีขั้นบันไดจะปฏิเสธสมมติฐานว่างที่ไม่ขึ้นอยู่กับการจำนวนของสมมติฐานเพียงอย่างเดียว แต่จะขึ้นอยู่กับผลลัพธ์ของการทดสอบสมมติฐาน อื่นๆ ด้วย (Dudoit et al, 2003)

ในกรณีที่เราไม่ทราบการแจกแจงร่วมของสถิติทดสอบ เราจะใช้วิธีการสุ่มตัวอย่างซ้ำ (resampling) เช่น วิธีการเรียงสับเปลี่ยน (permutation) และบูทสเตรป ซึ่งสามารถใช้ประมาณค่าตัวประมาณทั้งแบบปรับค่าพีและไม่ปรับค่าพี แต่ต้องระมัดระวังในเรื่องข้อตกลงเบื้องต้นของพารามิเตอร์ที่เกี่ยวข้องกับการแจกแจงร่วมของสถิติทดสอบด้วย

ในปี 1993 เวสต์ฟอลและยัง (Westfall and Young) ได้เสนอการทดสอบสมมติฐานของเปรียบเทียบพหุคูณ โดยเปรียบเทียบค่า p-value ที่น้อยที่สุดกับระดับนัยสำคัญ α โดยที่ค่า p-value ที่น้อยที่สุดนี้ หาได้จาก $p = P(\text{MinP} \leq \min p)$ โดยที่ MinP คือค่า p-value ที่น้อยที่สุดที่ได้จากวิธีบูทสเตรป และสำหรับ $\min p$ คือค่า p-value ที่น้อยที่สุดในบรรดาค่า p-value ของสมมติฐานการทดสอบที่มีตัวแปรหรือทรีตเมนต์ที่ k โดยปกติแล้วเราไม่ทราบการแจกแจงของ MinP แต่เราจะประมาณค่า MinP ได้อย่างง่ายด้วยวิธีการบูทสเตรป เวสต์ฟอลและยังได้เสนอว่าค่า p-value ของสมมติฐานอินเตอร์เซกชันอาจหาได้จากค่า p-value ที่น้อยที่สุด

ในบรรดาค่า p-value ของสมมติฐานเชิงเดียว รวมถึงการเสนอวิธีการทดสอบสมมติฐาน 2 วิธี คือ วิธีเวสต์ฟอล-ยัง บุทสเตรป ที่มีการสุ่มตัวอย่างแบบคืนที่ และวิธีการเรียงสับเปลี่ยนที่แท้จริง ที่มีการสุ่มตัวอย่างแบบไม่คืนที่

โปรแกรม SAS® จะใช้คำสั่ง PROC MULTEST ที่ให้ค่าที่ถูกต้องและเป็นการทดสอบแบบปิด โดยเฉพาะอย่างยิ่งการเปรียบเทียบหลายกลุ่มเพื่อเปรียบเทียบกับทรีตเมนต์ควบคุม ซึ่งในคำสั่ง PROC MULTEST จะทำการปรับค่าพี (adjusted P-value) โดยใช้หลักการของ Bonferroni, Sidak, Bootstrap resampling และ Permutation resampling ทั้งวิธีการชั้นเดียว และวิธีการสเตปดาวน์ โดยการจะทำการทดสอบในคำสั่ง PROC MULTEST ใช้หลักในการพิจารณาปัญหา ดังนี้

- 1) ข้อตกลงของตัวแบบสถิติ
- 2) วัตถุประสงค์ของการเปรียบเทียบหรือการทดสอบ
- 3) สมาชิกในเซตที่เราต้องการเปรียบเทียบ เช่น เปรียบเทียบทุกคู่ในการวิเคราะห์ความแปรปรวน (ANOVA), เปรียบเทียบทุกคู่ทรีตเมนต์กับทรีตเมนต์ควบคุม หรือการเปรียบเทียบพหุคูณทุกคู่กับทรีตเมนต์ที่ดีที่สุด

ในปี 2001 สมิธ (Smith) และ เทย์เลอร์ (Taylor) ได้กล่าวถึงวิธีดิเพนเดนท์บุทสเตรป (Dependent Bootstrap Procedure) และคุณสมบัติที่สำคัญ ต่อมาพวกเขาได้พิจารณาหาช่วงความเชื่อมั่นโดยใช้วิธีการวิธีดิเพนเดนท์บุทสเตรป พบว่า ค่าที่ครอบคลุมช่วงความเชื่อมั่นที่ได้ใกล้เคียงกับวิธีดั้งเดิมของวิธีอินดิเพนเดนท์บุทสเตรป (traditional (independent) bootstrap) และทฤษฎีช่วงความเชื่อมั่นปกติ ซึ่งมีค่าครอบคลุมช่วงความเชื่อมั่นสั้นกว่าเล็กน้อย

ในปี 2551 ยุทธสิทธิ์ (Yuttasit) ได้ทำการศึกษาวิธีการของ เวสต์ฟอลและยัง (Westfall and Young) โดยพิจารณาวิธีดิเพนเดนท์ บุทสเตรป มิน พี (dependent bootstrap min p) ในการเปรียบเทียบค่าเฉลี่ยกรณีที่มีทรีตเมนต์ควบคุม โดยทำการศึกษาวิธีสเตปดาวน์ดิเพนเดนท์ บุทสเตรป มิน พี เปรียบเทียบกับสถิติทดสอบของดันเนตต์ ในกรณีที่ข้อมูลตัวอย่างมีการแจกแจงปกติและมีความแปรปรวนเท่ากันในทุกทรีตเมนต์ ซึ่งผลการศึกษาที่ได้ คือ วิธีสเตปดาวน์ดิเพนเดนท์ บุทสเตรป มิน พี มีประสิทธิภาพใกล้เคียงกับสถิติทดสอบของดันเนตต์ (Komonnirarn, Y., 2008)