

บทที่ 2

ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

ในบทนี้จะกล่าวถึงทฤษฎีและงานวิจัยที่ใช้ในการศึกษาประกอบการทำวิทยานิพนธ์ โดยส่วนแรกเป็นการอธิบายทฤษฎีทางสถิติที่ใช้ในงานวิจัย ส่วนถัดมาเป็นการนำเสนองานวิจัยที่เกี่ยวข้อง และตัวอย่างงานวิจัยที่ผ่านมา รายละเอียดดังต่อไปนี้

2.1 การวัดค่าปัจจัยและทฤษฎีทางสถิติที่ใช้ในงานวิจัย

2.1.1 การวัดค่าปัจจัยที่ใช้กำหนดคุณภาพการให้บริการ

โดยทั่วไปแล้วปัจจัยที่นำมาใช้กำหนดคุณภาพการให้บริการของเว็บไซต์สามารถแบ่งออกได้ 2 ประเภท คือปัจจัยที่มีค่าแน่นอน (Certain Attributes) เป็นปัจจัยที่สามารถกำหนดค่าได้ มีค่าคงที่ไม่เปลี่ยนแปลง เช่น ค่าใช้บริการ ส่วนอีกประเภทคือ ปัจจัยที่มีค่าไม่แน่นอน (Uncertain Attributes) ค่าของปัจจัยได้จากการใช้บริการ ไม่สามารถกำหนดค่าที่แน่นอนได้ และค่าของปัจจัยขึ้นอยู่กับสภาพแวดล้อมในการทำงาน เช่น เวลาตอบสนอง โดยการอธิบายค่าปัจจัยที่มีค่าไม่แน่นอนเหล่านี้สามารถทำได้ 3 วิธี (Cheng Wan and Hongbing Wang, 2007) ดังนี้

1. การใช้สูตรคำนวณ (Function)
2. การใช้ขอบเขตบนและขอบเขตล่างของช่วงข้อมูลที่เป็นไปได้
3. การใช้ข้อมูลจากการให้บริการในอดีต (Historical Data)

2.1.2 ทฤษฎีทางสถิติ

2.1.2.1 ประเภทของสถิติ

สถิติ (Statistics) เป็นสถิติศาสตร์หรือวิชาสถิติที่ว่าด้วยหลักการและระเบียบวิธีอันได้แก่ การเก็บรวบรวม จัดระเบียบ สรุป นำเสนอ และวิเคราะห์ข้อมูล รวมทั้งทำข้อสรุป และตัดสินใจอย่างมีเหตุผลบนพื้นฐานของข้อมูลที่ได้จากการวิเคราะห์แล้ว (บุญธรรม กิจปรีดาบริสุทธิ์, 2549) สถิติแบ่งออกเป็น 2 ประเภท ได้แก่

1. สถิติพรรณนา (Descriptive Statistics) เป็นสถิติที่ใช้อธิบายคุณสมบัติต่าง ๆ ของสิ่งที่ต้องการศึกษาในกลุ่มใด เช่น สัดส่วน ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน เป็นต้น

2. สถิติอ้างอิง (Inferential Statistics) เป็นสถิติที่วิเคราะห์คุณสมบัติข้อมูลกลุ่มตัวอย่าง แล้วอ้างอิงไปยังกลุ่มประชากรเป้าหมาย ซึ่งมีวัตถุประสงค์เพื่อพยากรณ์ (Predict) และประมาณค่า (Estimate) คุณสมบัติของประชากรเป้าหมาย

2.1.2.2 การวิเคราะห์ข้อมูล (กัลยา วานิชย์บัญชา, 2545)

การวิเคราะห์ข้อมูลขั้นต้น หรือเรียกว่าสถิติพรรณนา เป็นการสรุปลักษณะเบื้องต้นของข้อมูลที่เกิดขึ้นรวบรวมมาได้ ผู้ที่ทำกรวิเคราะห์ข้อมูลขั้นต้นอาจจะเป็นผู้ที่ไม่มีความรู้สถิติมาก่อนก็ได้ การวิเคราะห์ข้อมูลขั้นต้นอาจพิจารณาในรูปแบบของการแจกแจงความถี่ การหาสัดส่วนหรือร้อยละ การวัดแนวโน้มเข้าสู่ส่วนกลาง เช่น การหาค่าเฉลี่ย ค่ามัธยฐาน และค่าฐานนิยม การวัดการกระจายของข้อมูล เช่น พิสัย ค่าแปรปรวน ค่าเบี่ยงเบนมาตรฐาน ฯลฯ

การสร้างตารางแจกแจงความถี่และกราฟ เป็นการนำเสนอข้อมูลในรูปตารางสำหรับตารางที่ใช้สรุป หรือนำเสนอข้อมูลนั้นอาจจะเป็นตารางแบบทางเดียว แบบสองทาง และแบบหลายทาง ตารางเหล่านี้เป็นตารางแจกแจงความถี่โดยการนำข้อมูลดิบซึ่งเป็นข้อมูลที่ถูกรวบรวม ในกรณีที่มีข้อมูลดิบเป็นจำนวนมาก เราจะนำข้อมูลดิบนั้นมาแบ่งออกเป็นกลุ่ม ๆ โดยอาจจะจำแนกตามลักษณะต่าง ๆ หรือจัดให้ข้อมูลที่มีค่าใกล้เคียงกัน หรือซ้ำกันอยู่ด้วยกัน เพื่อให้สรุปลักษณะของข้อมูลหรือทำการวิเคราะห์ข้อมูลได้ง่ายขึ้น การแจกแจงความถี่ทำได้ 2 แบบ คือ

1. การแจกแจงความถี่ของลักษณะที่สนใจที่เป็นไปได้ทั้งหมด หรือการแจกแจงความถี่ของข้อมูลเชิงคุณภาพ

2. การแจกแจงความถี่สำหรับค่าในแต่ละช่วงของลักษณะที่สนใจ หรือการแจกแจงความถี่ของข้อมูลเชิงปริมาณ

การสร้างตารางแจกแจงความถี่ทำได้ 2 ลักษณะ

1. กรณีที่กำหนดจำนวนชั้นหรือจำนวนกลุ่มไว้ล่วงหน้า

การตัดสินใจว่าควรให้มีข้อมูลกี่ชั้นหรือกี่กลุ่มนั้นพิจารณาจากจำนวนข้อมูลที่มี ถ้าจำนวนข้อมูลมีจำนวนมาก ควรมีจำนวนชั้นมาก อย่างไรก็ตามโดยทั่วไปตารางแจกแจงความถี่ควรมีอย่างน้อย 5 ชั้นแต่ไม่มากกว่า 15 ชั้น กรณีที่ข้อมูลมีการกระจายมาก (ค่าแตกต่างกัน

มาก) ควรกำหนดให้มีจำนวนชั้นน้อย ๆ เพื่อป้องกันไม่ให้ชั้นที่มีความถี่เป็นศูนย์ เนื่องจากไม่มีข้อมูลค่าใดตกอยู่ในชั้นนั้น ๆ เลย

กรณีที่ ไม่ทราบว่าจะกำหนดจำนวนชั้นเป็นเท่าใด อาจจะใช้สูตรจำนวนชั้นดังนี้

$$K = \text{จำนวนชั้น} = 1 + 3.3 \log N \text{ เมื่อ } N = \text{จำนวนข้อมูล} \quad (2.1)$$

หลังจากที่ได้จำนวนชั้นแล้ว จะคำนวณค่าดังต่อไปนี้

- **หาค่าพิสัยของข้อมูล** โดยที่ค่าพิสัยคือผลต่างระหว่างข้อมูลที่มีค่ามากที่สุดกับข้อมูลที่มีค่าน้อยสุด ดังนี้

$$\text{พิสัย (Range)} = \text{ค่าสูงสุด} - \text{ค่าต่ำสุด} \quad (2.2)$$

- **คำนวณหาความกว้างของชั้น (Class Interval : I)**

$$I = \frac{\text{ความกว้างของชั้น}}{\text{จำนวนชั้น}} \quad (2.3)$$

ถ้า I เป็นเลขไม่ลงตัว จะปัดขึ้นให้เป็นจำนวนเต็ม (ไม่ว่าเศษจะมีค่าต่ำกว่าหรือมากกว่า 0.5) โดยทั่วไปมักกำหนดให้ความกว้างของแต่ละชั้นเท่ากันหมด แต่ในทางปฏิบัติ บางครั้งอาจจะให้ความกว้างของแต่ละชั้นไม่เท่ากัน หรืออาจกำหนดให้เป็นชั้นเปิดก็ได้ หรืออาจกำหนดให้ความกว้างของชั้นเป็นค่าที่ทำให้ค่ากลาง (Mid Point) มีค่าเท่ากับค่าจริงของข้อมูล

โดยที่

$$\text{ค่ากลาง} = (\text{ขอบเขตจำกัดบน} + \text{ขอบเขตจำกัดล่าง}) / 2$$

$$\text{หรือ} \quad \text{ค่ากลาง} = (\text{ขีดจำกัดบน} + \text{ขีดจำกัดล่าง}) / 2 \quad (2.4)$$

- **ค่านวนหาขีดจำกัดชั้น (Class Limit)**

โดยจะกำหนดให้ขีดจำกัดล่างของชั้นแรก (ชั้นที่มีค่าต่ำสุด) ครอบคลุมข้อมูลที่มีค่าต่ำสุดและให้ขีดจำกัดบนของชั้นสุดท้าย (ชั้นที่มีค่าสูงสุด) ครอบคลุมข้อมูลที่มีค่าสูงสุด

- **ค่านวนหาขอบเขตจำกัดชั้น (Class Boundaries)**

การหาขอบเขตชั้นนั้นจะกำหนดให้ขอบเขตชั้นมีจำนวนหลักหลังจุดทศนิยมมากกว่าค่าของข้อมูลจริงอยู่ 1 หลักเสมอ เช่น ถ้าข้อมูลจริงเป็นเลขจำนวนเต็ม ขอบเขตจำกัดชั้นจะมีจุดทศนิยม 1 หลัก ในทางปฏิบัติ เราสามารถหาค่าขอบเขตของชั้นได้ดังนี้

$$\text{ขอบเขตจำกัดชั้น} = (\text{ขีดจำกัดบนของชั้น} + \text{ขีดจำกัดล่างของชั้นถัดไป})/2 \quad (2.5)$$

- **นับจำนวนค่าของข้อมูล (ความถี่) ในแต่ละชั้น**

หลังจากสร้างขอบเขตจำกัดชั้นแล้ว จึงตรวจสอบว่าข้อมูลค่าใดอยู่ในชั้นใดบ้าง แล้วนับข้อมูลในแต่ละชั้น เรียกว่า ความถี่ของชั้น

2. กรณีที่กำหนดความกว้างของชั้น

ขั้นตอนเหมือนการกำหนดจำนวนชั้น ยกเว้น การคำนวณหาจำนวนชั้นโดยที่กำหนดความกว้างแต่ละชั้นไว้ล่วงหน้า

$$\text{จำนวนชั้น} = \frac{\text{พิสัย}}{\text{ความกว้างของแต่ละชั้นที่กำหนด}} \quad (2.6)$$

2.1.2.3 การวิเคราะห์ตัวแทนของกลุ่ม

การวัดแนวโน้มเข้าสู่ศูนย์กลาง (Measure of Central Tendency) เป็นการอธิบายลักษณะข้อมูลที่สะดวกและเข้าใจง่าย วิธีที่นิยมและมีการใช้กันอย่างกว้างขวาง ได้แก่ ค่าเฉลี่ย (Mean) มัธยฐาน (Median) และฐานนิยม (Mode) ซึ่งแต่ละวิธีมีข้อดีและข้อเสียที่แตกต่างกัน (บุญธรรม กิจปรีดาบริสุทธิ์, 2549, น. 40)

ค่าเฉลี่ย (Mean) หรือตัวกลางเลขคณิต (Arithmetic Mean) เป็นผลเฉลี่ยข้อมูลทุกตัวในตัวแทนชุดนั้น เหมาะสำหรับข้อมูลที่มีลักษณะการแจกแจงแบบโค้งปกติ (Normal Distribution)

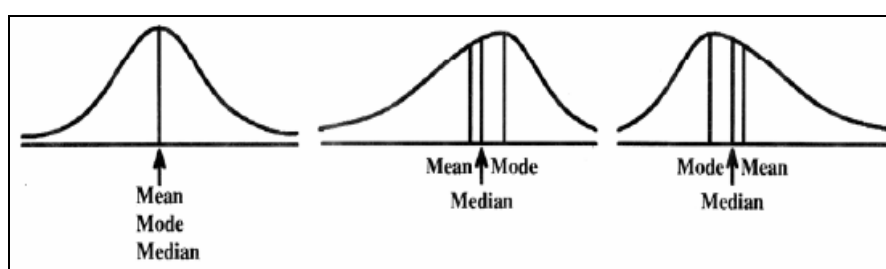
มัธยฐาน (Median) เป็นค่าของข้อมูลที่อยู่ตรงกลางกลุ่มข้อมูลที่มีการเรียงลำดับแล้ว

ในการวัดแนวโน้มเข้าสู่ศูนย์กลาง มัธยฐานจะเหมาะสมสำหรับข้อมูลที่มีลักษณะการแจกแจงแบบเบ้ ทั้งเบ้ซ้ายและเบ้ขวาได้ดีกว่าตัวอื่น

ฐานนิยม (Mode) เป็นค่าของข้อมูลในกลุ่มที่มีความถี่หรือมีค่าซ้ำกันมากที่สุด ควรใช้เมื่อข้อมูลมีลักษณะการแจกแจงแบบโด่งเป็น 2 แห่ง

ภาพที่ 2.1

แสดงตำแหน่งของค่าเฉลี่ย มัธยฐาน และฐานนิยมในการแจกแจงแบบต่าง ๆ



ที่มา : <http://www.watpon.com/Elearning/stat7.htm>

การอธิบายลักษณะข้อมูลด้วยการวัดแนวโน้มเข้าสู่ศูนย์กลางเพียงอย่างเดียวยังสรุปลักษณะข้อมูลที่นำมาใช้ได้ไม่สมบูรณ์ จะต้องบอกการกระจาย (Dispersion) ซึ่งการวัดการกระจาย (Measures of Dispersion or Variability) เป็นการบอกความแตกต่างกันภายในกลุ่มข้อมูล และมีสถิติที่สำคัญ (บุญธรรม กิจปริดาภิสุทธิ, 2549, น. 56) ดังนี้

พิสัย (Range) หมายถึงผลต่างของค่าสูงสุดและค่าต่ำสุดของข้อมูลชุดหนึ่ง บอกรายละเอียดของข้อมูลได้ไม่มากนัก เนื่องจากคำนวณจากข้อมูลเพียงบางส่วนเท่านั้น

ความแปรปรวน (Variance) เป็นค่าเฉลี่ยของผลรวมของผลต่างกำลังสอง ส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation) คือค่ารากที่สองของความแปรปรวน

2.1.2.4 การทดสอบสมมติฐาน (Hypothesis Testing)

การทดสอบสมมติฐาน (Hypothesis Testing) เป็นวิธีการหนึ่งในสถิติอ้างอิงเพื่อพิสูจน์ว่าค่าพารามิเตอร์ (Parameter) ในประชากรเป้าหมายแตกต่าง หรือสัมพันธ์เกี่ยวข้องกับค่าอุดมคติ ค่าทฤษฎี หรือค่าที่คาดหวังหรือไม่ หรือมีความแตกต่างกันระหว่างประชากรเป้าหมายหรือไม่ (บุญธรรม กิจปริดาภิสุทธิ, 2549, น. 81)

การทดสอบสมมติฐานจะเกี่ยวข้องกับประชากร (Population) และกลุ่มตัวอย่าง (Sample) โดยมีหน่วยตัวอย่าง (Sample Unit) เป็นข้อมูลที่สนใจศึกษา และประชากรเป็นกลุ่มของหน่วยตัวอย่างทั้งหมดที่ต้องการศึกษา ซึ่งมีค่าพารามิเตอร์ เป็นคุณลักษณะของประชากรที่คำนวณได้ด้วยวิธีทางสถิติจากข้อมูลประชากรทั้งหมด เช่น ค่าเฉลี่ยของประชากร ซึ่งจะแทนด้วยสัญลักษณ์ μ และความแปรปรวนของประชากร แทนด้วยสัญลักษณ์ σ^2 เป็นต้น ส่วนตัวอย่างเป็นกลุ่มย่อยของประชากรที่ถูกเลือกมาเพื่อทำการศึกษา โดยมีค่าสถิติเป็นคุณลักษณะของตัวอย่าง เช่น ค่าเฉลี่ยของตัวอย่าง ซึ่งจะแทนด้วยสัญลักษณ์ $\bar{x}$ และความแปรปรวนของตัวอย่าง แทนด้วยสัญลักษณ์ S^2 หรือ SD^2 เป็นต้น ซึ่งค่าสถิติเหล่านี้จะสอดคล้องกับค่าพารามิเตอร์ของประชากร และใช้เป็นค่าประมาณ (Estimator) ของค่าพารามิเตอร์โดยอาศัยทฤษฎีความสัมพันธ์ระหว่างค่าสถิติกับค่าพารามิเตอร์ซึ่งอยู่บนพื้นฐานของความน่าจะเป็น (อุมาพร จันทกร, 2542, น. 9) โดยขั้นตอนการทดสอบสมมติฐานแบ่งได้ 6 ขั้นตอนดังนี้ (บุญธรรม กิจปรีดาบริสุทธิ์, 2549, น.112-114)

1. การตั้งสมมติฐานเป็นการตั้งสมมติฐานทางสถิติเพื่อใช้ในการตัดสินใจเกี่ยวกับลักษณะประชากร ซึ่งจะแบ่งออกเป็น 2 ประเภท ได้แก่ สมมติฐานหลัก (Null Hypothesis) เป็นข้อความหรือสมการที่ใช้ทดสอบค่าพารามิเตอร์ แทนด้วยสัญลักษณ์ H_0 และ สมมติฐานเลือก (Alternative Hypothesis) เป็นข้อความหรือสมการที่ตั้งขึ้นเพื่อให้เป็นทางเลือกของโอกาสที่จะเกิดขึ้น เขียนแทนด้วยสัญลักษณ์ H_a

2. กำหนดระดับนัยสำคัญ (Level of Significance) เป็นการกำหนดโอกาสของการเกิดความคลาดเคลื่อนชนิดที่ 1 (Type I Error) ซึ่งความคลาดเคลื่อนชนิดที่ 1 คือ ความคลาดเคลื่อนที่เกิดจากการปฏิเสธสมมติฐานหลัก (H_0) ทั้งที่สมมติฐานหลักเป็นจริง ส่วนสัญลักษณ์ที่ใช้แทนระดับนัยสำคัญคือ α และปกติจะกำหนดไว้ที่ 0.05 หรือ 0.01

3. คำนวณค่าสถิติ เป็นการคำนวณค่าสถิติที่ใช้ทดสอบ ตัวอย่างเช่น ทดสอบค่าเฉลี่ยของตัวอย่าง 2 กลุ่มที่เป็นอิสระต่อกันโดยใช้การทดสอบที่ใช้พารามิเตอร์ จะใช้การทดสอบ-ที (T-Test) ค่าสถิติที่จะถูกคำนวณจากกลุ่มตัวอย่างด้วยสูตรของการทดสอบ-ที

4. หาค่าวิกฤตหรือค่าพี (Critical Value or P-Value) ค่าวิกฤตเป็นค่าที่แสดงขอบเขตบริเวณวิกฤต (Critical Region) ซึ่งค่าวิกฤตจะถูกกำหนดโดยค่าระดับนัยสำคัญ (α) และชนิดการทดสอบ ส่วนค่าพี (P-Value) คือ ระดับนัยสำคัญที่แท้จริง เป็นค่าความน่าจะเป็นหรือโอกาสผิดพลาดที่เกิดจากการปฏิเสธสมมติฐานหลัก (H_0) ทั้งที่สมมติฐานหลักเป็นจริง

5. การตัดสินใจทางสถิติ ในกรณีใช้ค่าวิกฤติ ถ้าค่าสถิติที่คำนวณได้ตกอยู่ในขอบเขตวิกฤติ จะถือว่าการทดสอบนั้นมีนัยสำคัญ (Significance) นั่นคือจะทำการปฏิเสธสมมติฐานหลัก และยอมรับสมมติฐานเล็กลง ส่วนในกรณีที่ใช้ค่าพี ถ้าทำการทดสอบแบบทางเดียว เมื่อค่า $P \leq \alpha$ จะปฏิเสธสมมติฐานหลัก (Reject H_0) ในทางตรงข้ามถ้า $P > \alpha$ จะยอมรับสมมติฐานหลัก (Accept H_0)

6. สรุปผล เป็นการสรุปผลการทดสอบสมมติฐานจากผลการตัดสินใจทางสถิติ

2.1.2.5 สถิติทดสอบ

การทดสอบสมมติฐานแบ่งการทดสอบออกเป็น 2 ประเภท ได้แก่ การทดสอบที่ใช้พารามิเตอร์ (Parametric Test) และการทดสอบที่ไม่ใช้พารามิเตอร์ (Non-Parametric Test) (บุญธรรม กิจปรีดาบริสุทธิ์, 2549, น. 81)

การทดสอบที่ใช้พารามิเตอร์เป็นการทดสอบที่มีข้อตกลงเบื้องต้นจำนวนมากและเข้มงวด ในการทดสอบความแตกต่างกันของประชากร 2 กลุ่มจะใช้การทดสอบ-ที ซึ่งมีข้อกำหนดเบื้องต้นคือ กลุ่มตัวอย่างต้องมาจากประชากรที่มีการแจกแจงแบบปกติ และในการทดสอบความแตกต่างกันของประชากรมากกว่า 2 กลุ่มจะใช้การวิเคราะห์ความแปรปรวนทางเดียว (One-way Analysis of Variance: Anova) ซึ่งจะมีข้อกำหนดเบื้องต้นเพิ่มขึ้น คือ ตัวอย่างต้องมาจากการประชากรที่มีความแปรปรวนเท่ากัน ซึ่งถ้าละเมิดข้อกำหนดเบื้องต้นผลที่ได้จะไม่เป็นที่ไว้วางใจ (Valid) และเชื่อถือได้

การทดสอบที่ไม่ใช้พารามิเตอร์จะเป็นการทดสอบที่ไม่มีข้อกำหนดเบื้องต้นเกี่ยวกับการแจกแจงของประชากรที่ตัวอย่างถูกสุ่มมา หรือเรียกว่าเป็นวิธีการแบบ Distributions-Free ซึ่งจะใช้แทนการทดสอบที่ใช้พารามิเตอร์เมื่อข้อมูลไม่เป็นไปตามข้อกำหนดเบื้องต้น (อุมาพร จันทศร, 2542, น. 2) และการทดสอบที่ไม่ใช้พารามิเตอร์จะมีคุณสมบัติที่เรียกว่าแกร่ง (Robust) คือ จะให้ผลการทดสอบที่ไว้วางใจ (Valid) แม้ว่าจะละเมิดข้อกำหนดเบื้องต้นหนึ่งข้อหรือมากกว่า (อุมาพร จันทศร, 2542, น. 24)

สำหรับปัญหาตำแหน่งที่ตั้งของตัวอย่าง 1 กลุ่มของการทดสอบที่ไม่ใช้พารามิเตอร์นั้น สามารถใช้การทดสอบเครื่องหมาย (Signed Test) ในการประมาณค่าของข้อมูล ดังนี้

1. การทดสอบเครื่องหมาย (Signed Test) (สายชล สีนสมบูรณ์ทอง, 2552)

- ข้อมูล (Data)

ค่าสังเกตมี n ค่า คือ $Z_1, \dots, Z_n$ ที่ไม่มีการจับคู่

- ข้อสมมติ (Assumptions)

1. Z_i แต่ละตัวเป็นอิสระกันอย่างเด็ดขาด (Mutually Independent)

2. Z_i แต่ละตัวมาจากประชากรแบบต่อเนื่องเหมือนกันที่มีค่า θ

$$\text{คือ } P(Z_i > \theta) = \frac{1}{2}; i = 1, \dots, n$$

- การประมาณค่าแบบจุดและการประมาณค่าแบบช่วง θ (Point Estimation and Interval Estimation of θ)

ค่าประมาณของ θ คือ

$$\begin{aligned} \theta &= \text{median}\{Z_i, 1 \leq i \leq n\} \\ &= \begin{cases} Z^{(k)Z^{(k+1)}} + Z^{(k+1)} & ; n = 2k + 1 (n \text{ เป็นเลขคี่}) \\ \frac{2}{2} & ; n = 2k (n \text{ เป็นเลขคู่}) \end{cases} \end{aligned}$$

ช่วงความเชื่อมั่น $(1-\alpha)100\%$ ของ θ คือ

$$(\theta_L, \theta_U) = \left(Z^{(C_\alpha)}, Z^{(n+1-C_\alpha)} = Z^{\left(\frac{b_{\alpha,1}}{2,2}\right)} \right) \quad (2.7)$$

$$\text{โดยที่ } C_\alpha = n + 1 - b_{\frac{\alpha,1}{2,2}}$$

การประมาณค่าเมื่อตัวอย่างมีขนาดใหญ่ (Large-Sample Approximation)
เมื่อ n มีค่ามาก ประมาณค่า C_α จาก

$$C_\alpha = \frac{n}{2} - Z_{\frac{\alpha}{2}} \sqrt{\frac{n}{4}}$$

โดยทั่วไปค่าด้านขวามือของสมการนี้ไม่เป็นจำนวนเต็ม จึงให้ C_α มีค่าเท่ากับจำนวนเต็มที่สูงที่สุด ซึ่งน้อยกว่าหรือเท่ากับค่าด้านขวามือของสมการนี้

สำหรับการทดสอบความแตกต่างของค่าแนวโน้มเข้าสู่ส่วนกลางระหว่างประชากรตั้งแต่ 2 กลุ่มขึ้นไปที่เป็นอิสระกัน ในสถิติที่ไม่ใช้พารามิเตอร์นั้นจะใช้ค่ามัธยฐานเป็นค่าสถิติในการทดสอบ โดยจะใช้การทดสอบของแมนทิวินี (Mann-Whitney Test) ทำการทดสอบความแตกต่างของค่ามัธยฐานระหว่างประชากร 2 กลุ่มที่เป็นอิสระต่อกันแทนการทดสอบ-ที และการทดสอบของครัสคาลและวอลลิส (The Kruskal-Wallis One-Way Analysis of Variance by Ranks Test) ทำการทดสอบความแตกต่างของค่ามัธยฐานระหว่างประชากรมากกว่า 2 กลุ่มขึ้นไปที่เป็นอิสระต่อกันแทนการวิเคราะห์ความแปรปรวนทางเดียว (อุมาพร จันทศร, 2542, น. 2) โดยการทดสอบทั้งสองนี้เป็นการทดสอบแบบผลรวมอันดับที่ (Rank Sum Test) ซึ่งมีแนวคิดในการทดสอบคือ เมื่อนำข้อมูลที่ได้จากตัวอย่างของแต่ละประชากรมาเรียงตามค่าของข้อมูลจากน้อยไปมากและหาผลรวมของอันดับที่ของข้อมูลในแต่ละกลุ่มตัวอย่าง ในกรณีที่ค่ามัธยฐานของประชากรเท่ากัน ผลรวมของอันดับที่ของแต่ละกลุ่มตัวอย่างควรจะมีค่าใกล้เคียงกันหรือไม่แตกต่างกันมากในทางกลับกันถ้าผลรวมของอันดับที่ของแต่ละกลุ่มตัวอย่างมีค่าแตกต่างกันมาก ค่ามัธยฐานของประชากรที่นำมาเปรียบเทียบจะแตกต่างกัน อย่างไรก็ตามในกรณีที่ข้อมูลแต่ละกลุ่มตัวอย่างไม่เท่ากัน ค่าเฉลี่ยผลรวมอันดับที่จะถูกนำมาเปรียบเทียบแทนค่าผลรวมอันดับที่ ซึ่งต่อไปนี้จะกล่าวถึงรายละเอียดของการทดสอบทั้งสอง

1. การทดสอบของแมนทิวินี (Mann-Whitney Test)

การทดสอบของแมนทิวินีเป็นการทดสอบความแตกต่างของค่ามัธยฐานของประชากร 2 กลุ่มที่เป็นอิสระต่อกัน โดยมีสมมติฐานสำหรับการทดสอบแบบสองหาง ดังนี้

$$H_0 : M_1 = M_2$$

$$H_a : M_1 \neq M_2$$

โดยที่ M_1 และ M_2 แทนค่ามัธยฐานของประชากรที่ 1 และ 2 ตามลำดับ

วิธีการคำนวณค่าสถิติทดสอบจะนำข้อมูลทั้งสองกลุ่มมารวมกันให้เป็นชุดเดียวกัน จากนั้นนำข้อมูลมาเรียงลำดับจากน้อยไปมาก ในกรณีที่มีข้อมูลซ้ำกันจะใช้ลำดับที่เฉลี่ย จากนั้นแยกลำดับตามข้อมูลเป็น 2 กลุ่มเหมือนเดิม รวมลำดับที่ของแต่ละกลุ่ม และใช้กลุ่มที่มีจำนวนลำดับน้อยแทนค่าในสูตรคำนวณสำหรับการทดสอบ (บุญธรรม กิจปรีดาภิรุทธิ์, 2549, น.296) รายละเอียดดังต่อไปนี้

ถ้าขนาดกลุ่มตัวอย่างแต่ละกลุ่มน้อยกว่า 20 ใช้สูตรคำนวณดังนี้

$$T = S - \frac{n(n+1)}{2} \quad (2.8)$$

เมื่อ S = ผลรวมของลำดับที่ของกลุ่มตัวอย่างที่น้อยกว่า
 n = ขนาดของกลุ่มตัวอย่างในกลุ่มที่มีผลรวมของลำดับที่น้อยกว่า
 m = ขนาดของกลุ่มตัวอย่างในกลุ่มที่มีผลรวมของลำดับที่มากกว่า

ถ้าขนาดกลุ่มตัวอย่างของแต่ละกลุ่มหรือกลุ่มใดกลุ่มหนึ่งมากกว่า 20 ใช้สูตรคำนวณดังนี้

$$Z = \frac{T - \frac{mn}{2}}{\sqrt{\frac{nm(n+m+1)}{12}}} \quad (2.9)$$

การหาค่าวิกฤตและการสรุปผลการตัดสินใจแบ่งออกเป็น 2 กรณีตามสูตรคำนวณ ถ้าคำนวณด้วยสูตร (2.8) ค่าวิกฤต $W_{\frac{\alpha}{2}}$ หาได้จากตาราง Quantiles of the Mann-Whitney Test Statistic ที่ $\frac{\alpha}{2}$ และถ้าค่าสถิติทดสอบ $T < W_{\frac{\alpha}{2}}$ หรือ $T > nm - W_{\frac{\alpha}{2}}$ จะปฏิเสธสมมติฐานหลัก (Reject H_0) แต่ถ้าคำนวณด้วยสูตร (2.9) ค่าวิกฤต $Z_{\frac{\alpha}{2}}$ หาได้จากตาราง Z ที่ $\frac{\alpha}{2}$ ถ้าค่าสถิติทดสอบ $Z < -Z_{\frac{\alpha}{2}}$ หรือ $Z > Z_{\frac{\alpha}{2}}$ จะปฏิเสธสมมติฐานหลัก (Reject H_0)

2. การทดสอบของครัสคาลและวอลลิส (The Kruskal-Wallis One-Way Analysis of Variance By Ranks Test)

การทดสอบของครัสคาลและวอลลิส เป็นการทดสอบความแตกต่างของค่ามัธยฐานของประชากรมากกว่า 2 กลุ่มขึ้นไป โดยมีสมมติฐานสำหรับการทดสอบแบบสองหาง ดังนี้

$$H_0 : M_1 = M_2 = \dots = M_k$$

$$H_a : M_1 \neq M_2 \neq \dots \neq M_k$$

โดยที่ $M_1, M_2, \dots, M_k$ แทนค่ามัธยฐานของประชากรที่ $1, 2, \dots, k$ ตามลำดับ

วิธีการคำนวณค่าสถิติทดสอบ ข้อมูลทุกกลุ่มจะถูกนำมารวมกันแล้วจัดลำดับจากคะแนนต่ำสุดไปหามากที่สุด หลังจากนั้นทำการแยกลำดับของแต่ละกลุ่ม นำค่าไปแทนในสูตรซึ่งเป็นสูตรประมาณค่าความแปรปรวนของลำดับ ถ้ากลุ่มตัวอย่างมาจากประชากรต่างกัน ค่าสถิติ H จะมีค่ามาก ผลการทดสอบจะปฏิเสธสมมติฐานหลัก (บุญธรรม กิจปรีดาบริสุทธิ์, 2549, น.303)

- กรณีที่ข้อมูลไม่ซ้ำกันใช้สูตรคำนวณดังนี้

$$H = \left(\frac{12}{N(N+1)} \sum_{i=1}^k \frac{R_i^2}{n_i} \right) - 3(N+1) \quad (2.10)$$

- เมื่อ k = จำนวนกลุ่มของประชากรที่ทำการทดสอบ
 R_i = ผลรวมของลำดับที่ของกลุ่มตัวอย่างที่ i เมื่อ $i = 1, 2, \dots, k$
 n_i = ขนาดของกลุ่มตัวอย่างที่ i
 N = จำนวนรวมของขนาดของกลุ่มตัวอย่างทั้งหมด โดยที่

$$N = \sum_{i=1}^k n_i$$

- กรณีที่มีข้อมูลซ้ำกันใช้สูตรคำนวณดังนี้

$$H_c = \frac{H}{1 - \frac{\sum T}{N^3 - N}} \quad (2.11)$$

- เมื่อ $\sum T = \sum (t^3 - t)$
 t = จำนวนข้อมูลที่ซ้ำกันแต่ละตำแหน่ง
 H = ค่าสถิติที่คำนวณได้จากสูตร (2.11)

การหาค่าวิกฤตและการสรุปผลการตัดสินใจแบ่งออกเป็น 2 กรณี ตามขนาดตัวอย่างที่นำมาคำนวณ

ถ้าจำนวนกลุ่มของประชากรที่ทำการทดสอบ 3 กลุ่ม ($k=3$) และขนาดของกลุ่มตัวอย่างแต่ละกลุ่มน้อยหรือเท่ากับ 5 ($n_i \leq 5$) ค่าวิกฤต H_α หาได้จากตาราง Kruskal-Wallis ที่ α และถ้าค่าสถิติทดสอบ $H \geq H_\alpha$ จะปฏิเสธสมมติฐานหลัก (Reject H_0)

ถ้าขนาดของกลุ่มตัวอย่างกลุ่มใดกลุ่มหนึ่งมากกว่า 5 ($n_i > 5$) จะใช้ค่าวิกฤต χ^2_{α} ที่ α , $df = k-1$ จากตาราง χ^2 และถ้าค่าสถิติทดสอบ $H \geq \chi^2_{\alpha}$ จะปฏิเสธสมมติฐานหลัก (Reject H_0)

เมื่อต้องการเปรียบเทียบกลุ่มตัวอย่างที่ i และ j ว่าต่างกันหรือไม่ คำนวณได้จากสูตร (บุญธรรม กิจปรีดาบริสุทธิ์, 2549, น. 304) ดังนี้

$$\left| \bar{R}_i - \bar{R}_j \right| \geq t_{\frac{\alpha}{2}} \sqrt{S^2 \frac{(N-1-H)}{N-k} \left(\frac{1}{n_i} + \frac{1}{n_j} \right)} \quad (2.12)$$

ให้ $\bar{R}_i$ = ค่าเฉลี่ยของผลรวมลำดับที่ของกลุ่มตัวอย่างที่ i เมื่อ $\bar{R}_i = \frac{R_i}{n_i}$

$\bar{R}_j$ = ค่าเฉลี่ยของผลรวมลำดับที่ของกลุ่มตัวอย่างที่ j เมื่อ $\bar{R}_j = \frac{R_j}{n_j}$

R_i, R_j = ผลรวมลำดับกลุ่มตัวอย่างกลุ่ม i และ j ที่ทำการเปรียบเทียบ

n_i, n_j = ขนาดกลุ่มตัวอย่างกลุ่ม i และ j ที่ทำการเปรียบเทียบ

H = ค่าสถิติที่คำนวณได้จากสูตร (2.11)

$t_{\frac{\alpha}{2}}$ = ค่า t ในตาราง t ที่ $\frac{\alpha}{2}$, $df = N - k$

N = จำนวนรวมของขนาดของกลุ่มตัวอย่างทั้งหมด โดยที่ $N = \sum_{i=1}^k n_i$

$S^2 = \frac{N(N+1)}{12}$ แต่ถ้ามีข้อมูลซ้ำกัน S^2 จะถูกปรับใหม่ ดังนี้

$$S^2 = \frac{N(N+1)}{12} - \frac{\sum T}{12(N-1)}$$

ถ้าผลสมการทางซ้ายมากกว่า หรือเท่ากับทางขวา แสดงว่าตัวอย่างกลุ่มที่ i และ j แตกต่างกันอย่างมีนัยสำคัญ

2.2 งานวิจัยที่เกี่ยวข้อง

ส่วนนี้กล่าวถึงงานวิจัยที่เกี่ยวข้อง ตัวอย่างงานวิจัยที่เกี่ยวกับการคัดเลือกเซอวิสมาดำเนินการ และงานวิจัยที่เสนออัลกอริทึมในการประกอบรวมเว็บเซอวิสตามเงื่อนไขที่กำหนดดังนี้

1. Jorge Cardoso et al. (2002) เสนอแบบจำลองที่ใช้ในการทำนายคุณภาพการให้บริการ (Predictive QoS Model) แบบอัตโนมัติ ซึ่งกำหนดปัจจัยที่ใช้ในแบบจำลอง 4 ปัจจัย คือ เวลา (Time), ค่าใช้จ่าย (Cost) เช่น ค่าใช้จ่ายที่ใช้ในการจัดเตรียมทรัพยากรสำหรับการประมวลผลงาน (Setup Cost) เป็นต้น, ความน่าเชื่อถือ (Reliability) และความถูกต้องของการดำเนินการ (Fidelity) ซึ่งกำหนดให้การประมาณค่าคุณภาพการให้บริการใน 2 ขั้นตอนคือ ขั้นตอนการออกแบบ โดยผู้ออกแบบกำหนดค่าเริ่มต้นจากข้อมูลที่ได้จากการทดสอบการทำงาน ขั้นตอนที่สองคือคำนวณค่าคุณภาพการให้บริการใหม่ (Re-Computed) ณ รั้นไทม์ โดยใช้แบบจำลองคณิตศาสตร์ (Mathematical Model) คำนวณคุณภาพการให้บริการแต่ละปัจจัยของเวิร์คโฟลว์ วิธีการคือใช้เ็ชชดับเบิลยูอาอัลกอริทึม (SWR Algorithm) ในการลดรูป (Reduction) โครงสร้างของเวิร์คโฟลว์แบบต่าง ๆ ให้เป็นเวิร์คโฟลว์ที่มีโครงสร้างแบบลำดับ (Sequential) เนื่องจากเวิร์คโฟลว์สนับสนุนกฎในการลดทอน สามารถทำได้ง่าย และผลลัพธ์ของเวิร์คโฟลว์ที่ได้จากการลดทอนง่ายต่อการทำความเข้าใจ งานวิจัยของ Jorge Cardoso et al. ทำการทดสอบความถูกต้อง โดยการสร้างแบบจำลองและทำการทดลองอยู่บนพื้นฐานของข้อมูลจริง ซึ่งผลการทดลองสรุปได้ว่าแบบจำลองมีความเหมาะสมที่จะใช้ในการทำนายและวิเคราะห์คุณภาพการให้บริการของเวิร์คโฟลว์

2. Zeng et al. (2003) เสนอแบบจำลองคุณภาพที่ใช้ในการคัดเลือกเว็บเซอวิสเพื่อนำมาประกอบรวม โดยพิจารณาถึงข้อกำหนดของการประกอบรวม (Global Constraint) เมื่อสิ้นสุดการดำเนินการ ซึ่งได้เสนอวิธีการเลือกเซอวิสของทั้งกระบวนการทำงาน (Global Planning Approach) เปรียบเทียบกับวิธีการเลือกเซอวิสของแต่ละงาน โดยมีแนวคิดที่ว่า การเลือกเซอวิสแบบเดียวสำหรับแต่ละงาน เมื่อนำมาประกอบรวม คุณภาพของบริการที่ได้รับอาจจะไม่ตรงตามความต้องการ หรือข้อกำหนดที่วางไว้ ปัจจัยที่ใช้กำหนดคุณภาพการให้บริการที่พิจารณาคือ

ค่าบริการ (Execution Price) ระยะเวลาการบริการ (Execution Duration) ความน่าเชื่อถือสภาพพร้อมใช้งาน และค่าความนิยม (Reputation)

นอกจากนั้น เสนอวิธีการประกอบรวมเว็บเซอวิสที่ใช้โปรแกรมเชิงเส้น (LP- Linear Programming) ในการเลือกแผนประกอบรวม (Execution Plan) แทนการสร้างแผนการประกอบรวมออกมาทุกกรณีที่เป็นไปได้ ซึ่งเว็บเซอวิสในแผนประกอบรวมที่ให้ค่าคุณภาพสูงสุดจะถูกคัดเลือกมาดำเนินการ ในการวัดค่าคุณภาพของแผนประกอบรวมจะใช้สมการวัตถุประสงค์ (Objective Function) รวมค่าปัจจัยต่าง ๆ เข้าด้วยกันเพื่อคำนวณเป็นค่าคุณภาพของบริการ นอกจากนี้ได้นำความแปรปรวนเข้ามาเป็นปัจจัยหนึ่งในการพิจารณา ซึ่งผู้ให้บริการทำหน้าที่กำหนดระดับน้ำหนักความแปรปรวนของแต่ละปัจจัย

ผลการทดลองของ Zeng et al. พบว่า การประกอบรวมเว็บเซอวิสที่เลือกเซอวิสทั้งกระบวนการมาดำเนินการ ให้ผลดีกว่าการเลือกเซอวิสแบบเดี่ยวสำหรับแต่ละงานมาดำเนินการ จากที่กล่าวมานั้น แม้รายงานวิจัยของ Zeng et al. จะมีเทคนิคการคัดเลือกกลุ่มผู้ให้บริการคล้ายกับงานวิจัยฉบับนี้ แต่อย่างไรก็ตามงานวิจัยฉบับนี้ได้เพิ่มการตัดสินใจเปลี่ยนไปให้ผู้ให้บริการสำรองมาดำเนินการแทน ณ วันใหม่ ร่วมด้วย เมื่อการดำเนินการของผู้ให้บริการรายเดิมมีโอกาสทำให้เวิร์คโฟลว์ไม่ประสบความสำเร็จตามเวลาที่กำหนด

3. Tao Yu and Kwei-Jay Lin (2004) เสนออัลกอริทึมในการเลือกเซอวิสที่มีคุณภาพดีที่สุด มาประกอบรวมเว็บเซอวิส ภายใต้ข้อกำหนดของคุณภาพการให้บริการ (QoS Constraint) ซึ่งมุ่งเน้นเฉพาะปัจจัยด้านเวลาที่ผู้ใช้บริการได้รับจากการเรียกใช้งานเซอวิส (End-to-End Delay) และนำแนวคิดของ Multi-Choice Knapsack Problem เป็นสมการในการแก้ปัญหาปัจจัยด้านเวลาที่ผู้ใช้บริการได้รับจากการเรียกใช้งานเซอวิส

นอกจากนี้ปัจจัยด้านความหน่วงของระบบเครือข่าย (Network Delay) มาร่วมพิจารณาในการเลือกเซอวิส โดยการส่งผลลัพธ์ระหว่างเซอวิสได้กำหนดค่าความหน่วงของระบบเครือข่ายไว้ล่วงหน้า และมีค่าเท่ากันทุกเซอวิส จากนั้นเสนอ 3 อัลกอริทึมในการเลือกเซอวิส คือ Exhaustive Search Algorithm , Dynamic Programming Algorithm และ Pisinger's Algorithm โดย Tao Yu and Kwei-Jay Lin ได้สรุปผลการทดลองดังนี้ Exhaustive Search Algorithm เป็นวิธีการแบบตรงไปตรงมา โดยนำเซอวิสทั้งหมดมาคำนวณค่าคุณภาพ (Utility Function) แล้วนำมาเปรียบเทียบกัน วิธีนี้ทำให้ได้เซอวิสที่ดีที่สุด แต่เสียเวลาและสิ้นเปลืองหน่วยความจำ ดังนั้น วิธีนี้เหมาะกับการประกอบรวมเว็บเซอวิสที่มีกระบวนการ

ประกอบรวมและจำนวนผู้ให้บริการของแต่ละเซอวิสจำนวนน้อย กรณีที่มีจำนวนเซอวิสสำหรับการประกอบรวมและจำนวนผู้ให้บริการของแต่ละเซอวิสจำนวนมาก สรุปได้ว่า Pisinger's Algorithm ใช้เวลาในการเลือกน้อยกว่า และมีประสิทธิภาพกว่า Dynamic Programming Algorithm

4. Rainer Berbner et al. (2006) เสนอ Heuristic Algorithm โดยใช้ Mixed Integer Programming เป็นสมการสำหรับการประกอบรวม โดยการนำข้อมูลค่าคุณภาพมาจัดเรียง จากนั้นใช้ Backtracking Algorithm ในการคำนวณแผนประกอบรวมที่เหมาะสม นอกจากนั้นวิเคราะห์ปัจจัยที่ส่งผลกระทบต่อประสิทธิภาพ 3 ปัจจัยคือ จำนวนงาน จำนวนเว็บเซอวิสของแต่ละงาน และเปอร์เซ็นต์ความต้องการของเงื่อนไข จากนั้นทำการวิเคราะห์ประสิทธิภาพด้วยการเปรียบเทียบ Heuristic Algorithms กับ Linear Integer Programming ซึ่งสรุปผลได้ว่า Heuristic Algorithm มีประสิทธิภาพดีกว่าแม้จะมีการเพิ่มจำนวนงานและจำนวนเว็บเซอวิสของแต่ละงาน

5. Cheng Wan and Hongbing Wang (2007) เสนอแบบจำลองคุณภาพการให้บริการเพื่อคัดเลือกเว็บเซอวิส โดยงานวิจัยมุ่งเน้นไปที่ปัจจัยที่มีค่าไม่แน่นอน ซึ่งวัดค่าจากการใช้บริการในอดีตที่ผู้ให้บริการ โดยปัจจัยที่นำมาใช้วัดคุณภาพจะมีการทำข้อตกลงระหว่างผู้ให้บริการและผู้ให้บริการ สมการวัตถุประสงค์จะถูกกำหนดเป็นมาตรฐานที่ใช้รวมค่าของปัจจัยต่าง ๆ เข้าด้วยกันเพื่อใช้ในการกำหนดคุณภาพ โดยมีแบบจำลอง NPTQOS ในการเปรียบเทียบคุณภาพแบบจำลอง NPTQOS ใช้การทดสอบที่ไม่ใช่พารามิเตอร์ ซึ่งนำสถิติทดสอบแมนทิวเนียนยู (Mann-Whiney U) มาใช้เปรียบเทียบความแตกต่างระหว่างค่ากลางของข้อมูลคุณภาพการให้บริการของ 2 เว็บเซอวิส อัลกอริทึมในการเปรียบเทียบจะใช้สถิติทดสอบแมนทิวเนียนยู เปรียบเทียบเว็บเซอวิสทีละคู่ เพื่อจัดลำดับคุณภาพการให้บริการ ในกรณีที่ค่ากลางของคู่ใดไม่มีความแตกต่างกันอย่างมีนัยสำคัญ ใช้วิธีการเปรียบเทียบความโด่ง (Kurtosis) และความเบ้ (Skewness) ในการจัดลำดับเว็บเซอวิสคู่นั้น นอกจากนั้น Cheng Wan and Hongbing Wang กล่าวว่าวิธีการในการคำนวณค่าปัจจัยจากข้อมูลในอดีตใช้ต้นทุนน้อยกว่าวิธีการ Logistic Regression และ Multiple Linear Regression ดังนั้นแบบจำลองคุณภาพการให้บริการจึงเหมาะสมสำหรับการคัดเลือกเซอวิสที่มีค่าไม่แน่นอน

6. Wentao Zhang et al. (2007) เสนอการประกอบรวมเว็บเซอร์วิส ที่มีการคัดเลือกเว็บเซอร์วิสมาดำเนินการแบบไดนามิก โดยแบ่งกลุ่มปัญหาตามเงื่อนไขที่ใช้ในการคัดเลือก และนำปัญหามาแสดงให้อยู่ในรูปของแบบจำลองคณิตศาสตร์ (Mathematical Model) นอกจากนั้นเสนออัลกอริทึมที่ใช้ในการคัดเลือกเซอร์วิสมาดำเนินการแก้ปัญหาที่เหมาะสมกับแต่ละกลุ่ม ดังนี้

- การคัดเลือกเว็บเซอร์วิสที่ไม่มีเงื่อนไขในการคัดเลือก ใช้ Greedy Algorithm ในการแก้ปัญหา โดยทำการคำนวณค่าคุณภาพการให้บริการของทุกเซอร์วิสที่ใช้ในการดำเนินการงานเดียวกัน ซึ่งเว็บเซอร์วิสที่มีค่าคุณภาพมากที่สุดจะได้รับการคัดเลือก

- การคัดเลือกเว็บเซอร์วิสที่มี 1 เงื่อนไขในการคัดเลือก เช่น เงื่อนไขด้านเวลาตอบสนอง ใช้ Multi-Choice Knapsack Problem ซึ่งเป็นปัญหาแบบ NP-Complete วิธีการทั่วไปคือสร้างเส้นทางที่เป็นไปได้ทุกเส้นทาง แล้วเลือกเส้นทางที่มีคุณภาพการให้บริการดีที่สุด แต่วิธีการนี้เป็นการใช้ทรัพยากรอย่างสิ้นเปลือง และไม่เหมาะสมกับปัญหาขนาดใหญ่ จึงได้เสนอ Heuristic Algorithm ในการคำนวณหาเส้นทางที่ไม่ต้องดีที่สุด แต่มีคุณภาพใกล้เคียงและอยู่ภายใต้เงื่อนไขที่กำหนด

- การคัดเลือกเว็บเซอร์วิสที่มีหลายเงื่อนไขในการคัดเลือก ใช้ Multi-Dimension Multi-Choice Knapsack Problem ซึ่งเป็นปัญหาแบบ NP-Complete วิธีการนี้ใช้ Heuristic Algorithm เหมือนการคัดเลือกเว็บเซอร์วิสที่มี 1 เงื่อนไข แต่เพิ่มเติมในส่วนของการต้องการและต้องแปลงเงื่อนไขให้อยู่ในรูป 1 มิติ (Dimension)

ผลการทดลองของ Wentao Zhang et al. พบว่าอัลกอริทึมที่ใช้ในการคัดเลือกให้ประสิทธิภาพการทำงานดี แม้จะมีการเพิ่มจำนวนงานและจำนวนเว็บเซอร์วิสของแต่ละงาน และอัตราการเพิ่มของเวลาที่ใช้ในการคัดเลือกก็เป็นอัตราที่เหมาะสม

7. คุณลัดดาวัลย์ (Laddawan Kulnarattna and Songsakdi Rongviriyapanish) (2009) เสนอแบบจำลองการคัดเลือกผู้ให้บริการ โดยใช้ปัจจัยที่มีค่าไม่แน่นอนเป็นเกณฑ์ในการคัดเลือก 3 ปัจจัย คือ เวลาตอบสนอง สภาพพร้อมใช้งาน ความน่าเชื่อถือ และประยุกต์ใช้สถิติที่ไม่ใช่พารามิเตอร์ในการคัดเลือก โดยนำการทดสอบของแมนทีวียันย์ และการทดสอบของครัสคาลและวอลลิส ซึ่งมีสมมติฐานการทดสอบเป็นการเปรียบเทียบค่ามัธยฐานของปัจจัยที่มีค่าไม่แน่นอนระหว่างเว็บเซอร์วิสที่นำมาคัดเลือก ขั้นตอนการทดสอบคือนำกลุ่มของเว็บเซอร์วิสที่ให้บริการงานเดียวกัน นำมาจัดเรียงลำดับตามคุณภาพการให้บริการในแต่ละด้าน คือด้านเวลา

ตอบสนอง ด้านสภาพพร้อมใช้งาน และด้านความน่าเชื่อถือ ซึ่งค่าปัจจัยวัดจากข้อมูลการใช้บริการเว็บเซอร์วิสในอดีตที่ฝั่งผู้ให้บริการ จากนั้นนำมาคำนวณคุณภาพการให้บริการโดยคำนวณจากลำดับคุณภาพในแต่ละด้านของเว็บเซอร์วิสโดยใช้ฟังก์ชันถ่วงน้ำหนัก เว็บเซอร์วิสที่มีค่าคุณภาพการให้บริการสูงสุดจะเป็นเว็บเซอร์วิสที่ถูกเลือกมาใช้ในแบบจำลอง นอกจากนี้ วิเคราะห์การเปลี่ยนแปลงคุณภาพของเซอร์วิส ณ เวลาปัจจุบันซึ่งอาจมีการเปลี่ยนแปลงไปจากคุณภาพเซอร์วิส ณ ช่วงเวลาที่นำข้อมูลมาทำการทดสอบ ทำให้ลำดับคุณภาพของเว็บเซอร์วิสเปลี่ยนแปลงไป ดังนั้นจึงทดสอบการเปลี่ยนแปลงคุณภาพ โดยการทดสอบมุ่งเน้นทดสอบกับเซอร์วิสที่มีค่าลำดับคุณภาพเท่ากับ 1 ของทั้งสามปัจจัย เพราะคาดการณ์ว่าการเปลี่ยนแปลงของลำดับที่ 1 ส่งผลกระทบมากที่สุดต่อการคัดเลือก งานวิจัยของคุณลัดดาวัลย์วัดผลความแม่นยำของแบบจำลอง โดยเสนอในรูปแบบของเปอร์เซ็นต์ความแม่นยำในแต่ละด้านทั้ง 3 ปัจจัยที่นำมาพิจารณาผลการทดสอบพบว่า แบบจำลองมีความแม่นยำในการคัดเลือกเว็บเซอร์วิสด้านสภาพพร้อมใช้งาน และความน่าเชื่อถือ 98 เปอร์เซ็นต์ ส่วนเปอร์เซ็นต์ความแม่นยำในการเรียงลำดับด้านสภาพพร้อมใช้งาน และความน่าเชื่อถือมีค่า 73 เปอร์เซ็นต์ ในขณะที่ความแม่นยำด้านเวลาตอบสนองต้องทำการปรับปรุงเทคนิคเพิ่มเติม

8. คุณวารุณี (Warunee Srinaprom and Songsakdi Rongviriyapanish) (2009) เสนอการพัฒนาระบบเวิร์คโฟลว์ที่มีการนำเอเจนต์เข้ามาช่วยในการเชื่อมต่อกับแอปพลิเคชันอื่นแบบไดนามิก (Dynamic) เพื่อเพิ่มความยืดหยุ่น และรองรับต่อการเปลี่ยนแปลงของกระบวนการทางธุรกิจ นอกจากนี้ออกแบบให้ระบบจัดการเวิร์คโฟลว์มีการคัดเลือกผู้ให้บริการมาดำเนินการแบบอัตโนมัติ การคัดเลือกพิจารณาจากปัจจัยเชิงคุณภาพ 3 ด้าน ได้แก่ สภาพพร้อมให้บริการ ความน่าเชื่อถือ และเวลาตอบสนอง โดยค่าของปัจจัยทั้งสามด้านนี้ได้มาจากการดำเนินการในอดีตของผู้ให้บริการ ผลการทดลองพบว่าต้องปรับปรุงความแม่นยำการคัดเลือกผู้ให้บริการด้านเวลาตอบสนองเพิ่มเติม ในด้านประสิทธิภาพและค่าใช้จ่าย (Overhead) ของการนำแบบจำลองการคัดเลือกผู้ให้บริการมาใช้กับระบบจัดการเวิร์คโฟลว์พบว่าทำให้ความสำเร็จของเวิร์คโฟลว์เพิ่มขึ้น 10.21 เปอร์เซ็นต์ ในขณะที่โดยเฉลี่ยแล้วค่าใช้จ่ายที่เกิดจากการนำแบบจำลองการคัดเลือกผู้ให้บริการมาใช้เพิ่มขึ้นไม่ถึง 1 เปอร์เซ็นต์ของเวลาที่ใช้ในการประมวลผลเซอร์วิสนั้น ซึ่งสรุปได้ว่ามีความคุ้มค่าที่จะนำแบบจำลองการคัดเลือกผู้ให้บริการที่วัดค่าปัจจัยคุณภาพจากในอดีตมาใช้ในการคัดเลือกผู้ให้บริการของเวิร์คโฟลว์ แต่อย่างไรก็ตามงานวิจัยของคุณวารุณีไม่ได้นำปัจจัยข้อบังคับด้านเวลาตอบสนองของเวิร์คโฟลว์มาพิจารณา

9. Dieter Schuller et al. (2010) เสนอวิธีการคัดเลือกเซอวิสสำหรับเวิร์คโฟลว์ที่มีความซับซ้อน (Complex Workflow) กล่าวคือ เวิร์คโฟลว์ที่มีโครงสร้างการทำงานแบบลำดับแบบขนาน แบบมีเงื่อนไข และแบบทำซ้ำ โดยวิธีการคัดเลือกเริ่มต้นจากการกำหนดปัญหาให้อยู่ในรูปของสมการทางคณิตศาสตร์ เพื่อใช้คำนวณค่าปัจจัยคุณภาพ ได้แก่ เวลาที่ใช้ดำเนินการ (Execution Time) ค่าบริการ (Costs) ความน่าเชื่อถือ (Reliability) และปริมาณงานที่ทำได้ (Throughput) จากนั้นทำการปรับสมการให้อยู่ในรูปของสมการเชิงเส้น และใช้เทคนิค ILP (Integer Linear Programming) ในการคำนวณค่าคุณภาพของเซอวิส แต่เนื่องจากเทคนิค ILP จะใช้เวลาในการคำนวณมากเมื่อมีการเพิ่มจำนวนงาน และจำนวนของเซอวิสที่นำมาคัดเลือก Dieter Schuller et al. จึงได้เสนอเทคนิค MLIP (Mixed Integer Linear Programming) ในการคำนวณ ซึ่งสรุปได้ว่าเทคนิค MLIP สามารถทำงานได้อย่างมีประสิทธิภาพ แม้มีการเพิ่มจำนวนกิจกรรมและผู้ให้บริการ จากงานวิจัยของ *Dieter Schuller et al.* จะเห็นได้ว่านำสมการเชิงเส้นมาใช้ในการคัดเลือกเซอวิสเหมือนงานวิจัยฉบับนี้

10. Shao-chong Li et al. (2010) เสนอวิธีการคัดเลือกเว็บเซอวิสที่ตรงกับความต้องการของผู้ใช้ 2 วิธีการ คือการจัดลำดับคุณภาพของเว็บเซอวิสตามค่าที่ได้จากการคำนวณ (Static Ranking) และจัดลำดับคุณภาพของเซอวิสจากการเรียนรู้ข้อมูลจากในอดีต (Dynamic Ranking) กล่าวคือ

- Static Ranking กำหนดปัญหาให้อยู่ในรูปของสมการทางคณิตศาสตร์ให้สอดคล้องกับปัจจัยที่พิจารณาและความต้องการของผู้ใช้ กล่าวคือ นำปัจจัยที่พิจารณามาแสดงอยู่ในรูปของเวกเตอร์ (Vector) จากนั้นนำมาคำนวณค่าปัจจัยคุณภาพตามฟังก์ชันถ่วงน้ำหนักของแต่ละปัจจัยที่กำหนด โดยเว็บเซอวิสที่มีค่าปัจจัยคุณภาพจากการคำนวณมากที่สุด คือเว็บเซอวิสที่มีคุณภาพดีที่สุด

- Dynamic Ranking จัดลำดับคุณภาพโดยการปรับค่าฟังก์ชันถ่วงน้ำหนักที่เหมาะสมแบบอัตโนมัติ ซึ่งวิเคราะห์จากความสัมพันธ์ของค่าฟังก์ชันถ่วงน้ำหนักที่กำหนดจากในอดีต (Historical Data) และใช้ Back Propagation อัลกอริทึม ซึ่งเป็นหนึ่งในเทคนิคของนิวรอลเน็ตเวิร์ก (Artificial Neural Network) ในการจัดลำดับคุณภาพ

งานวิจัยของ Shao-chong Li et al. ทำการทดลองโดยสร้างต้นแบบเพื่อจำลองการทำงาน กำหนด 5 ปัจจัยที่พิจารณา คือ เวลาตอบสนอง สภาพพร้อมใช้งาน อัตราความล้มเหลว ค่าใช้จ่ายจากการดำเนินการ และความน่าเชื่อถือ โดยผลการทดลองพบว่า ความคลาดเคลื่อน

ระหว่าง 2 วิธีการมีค่าเป็น 1.23 เปอร์เซนต์ สรุปได้ว่าวิธีการคัดเลือกเว็บเซอร์วิสที่นำเสนอช่วยเพิ่มความแม่นยำในการคัดเลือกผู้ให้บริการ จากงานวิจัยของ Shao-chong Li et al. จะเห็นได้ว่าทำการกำหนดปัญหาให้อยู่ในรูปของสมการ และนำข้อมูลจากในอดีตมาใช้ในการคำนวณเหมือนงานวิจัยฉบับนี้ แต่อย่างไรก็ตามยังขาดการพิจารณาถึงคุณภาพของเว็บเซอร์วิสเมื่อนำมาดำเนินการร่วมกัน