

## มารู้จัก Nonparametric กันเถอะ

ศิริลักษณ์ สุวรรณวงศ์

คณะวิทยาศาสตร์ มหาวิทยาลัยมหิดล

นักวิจัยเชิงปริมาณไม่ว่าจะเป็นนักวิจัยหน้าใหม่ หรือหน้าเก่าแทบทุกท่านมักเกิดปัญหาการเลือกใช้สถิติวิเคราะห์ การเลือกใช้สถิติวิเคราะห์ที่นิยมสำหรับงานวิจัยที่มีวัตถุประสงค์เพื่อการเปรียบเทียบกรณีที่มี 2 กลุ่มที่เป็นอิสระ คือ Independent t-test ถ้าเป็น 2 กลุ่มที่ไม่เป็นอิสระ คือ Dependent t-test และเมื่อต้องการเปรียบเทียบมากกว่า 2 กลุ่ม จะเลือกใช้ ANOVA (Analysis of Variance) ซึ่งอาจเป็น One-way ANOVA หรือ Two-way ANOVA สถิติวิเคราะห์เหล่านี้เป็นสถิติพารามิเตอร์ นักวิจัยจะต้องมีความระมัดระวังเกี่ยวกับข้อตกลงเบื้องต้นของสถิติวิเคราะห์ที่เลือกใช้ ข้อตกลงเบื้องต้น (assumption) ที่สำคัญได้แก่ ข้อมูลต้องมีการแจกแจงปกติ (normal distribution) และมีลักษณะสุ่ม (random) ซึ่งนักวิจัยส่วนมากมักคิดว่าเขาได้ตัวอย่างด้วยวิธีการสุ่มตัวอย่างแบบต่างๆ ไม่ว่าจะเป็นการเลือกตัวอย่างสุ่มอย่างง่าย (simple random sampling) หรือการเลือกตัวอย่างสุ่มหลายขั้นตอน (multi-stage sampling) ดังนั้นกลุ่มตัวอย่างที่ได้จะมีลักษณะสุ่ม ซึ่งความคิดดังกล่าวไม่ถูกต้องเพราะวิธีการสุ่มตัวอย่างไม่ได้แปลว่าข้อมูลที่ท่านได้จะต้องมีลักษณะสุ่ม และความคิดของนักวิจัยบางท่านที่ว่าเมื่อสุ่มตัวอย่างมากกว่า 30 สามารถแปลความว่าข้อมูลมีการแจกแจงปกติ สิ่งนี้เป็นความคิดที่ผิดหลักสถิติอย่างรุนแรงซึ่งอาจจะส่งผลให้ผลการวิเคราะห์ไม่เกิดประโยชน์ หรือผลงานวิจัยผิดพลาดได้ ดังนั้นนักวิจัยควรต้องมีความใส่ใจที่จะต้องนำข้อมูลที่เก็บรวบรวมได้มาทดสอบข้อตกลงเบื้องต้นก่อนจะเลือกใช้สถิติวิเคราะห์ เพราะความน่าเชื่อถือของการวิจัยส่วนหนึ่งเกิดจากการสรุปอ้างอิงผลการวิจัย ดังนั้นการเลือกใช้สถิติวิเคราะห์ข้อมูลที่เหมาะสมจึงมีความสำคัญ ปัญหาที่จะเกิดขึ้นต่อมาคือถ้าข้อมูลไม่เป็นไปตามข้อตกลงเบื้องต้น แล้วเราจะใช้สถิติใดวิเคราะห์ได้ ผู้เขียนจึงขอแนะนำให้เรามารู้จัก สถิติไม่อิงพารามิเตอร์ (Nonparametric Statistics) พร้อมทั้งตัวอย่างการประยุกต์ใช้สถิติไม่อิงพารามิเตอร์

สถิติไม่อิงพารามิเตอร์ หรือบางท่านอาจเรียกว่าสถิติที่ไม่ใช่พารามิเตอร์ หรือสถิติไร้พารามิเตอร์ สถิติวิเคราะห์นี้เกิดขึ้นในปี ค.ศ. 1942 โดย Jacob Wolfowitz (ค.ศ. 1910 - ค.ศ. 1981) เป็นชาวโปแลนด์ได้ศึกษาสถิติประเภทใหม่โดยไม่สนใจลักษณะการแจกแจงของประชากร ซึ่งเขาเรียกสถิตินี้ว่า สถิติการแจกแจงอิสระ (Distribution-Free Statistics) และเป็นสถิติทดสอบสมมุติฐานที่ไม่เกี่ยวกับพารามิเตอร์ของประชากร และสถิติไม่อิงพารามิเตอร์ใช้ได้กับข้อมูลที่มีมาตราวัดระดับต่ำสุด คือ ระดับมาตราวัดนามบัญญัติ (Nominal scale) เมื่อกล่าวเช่นนี้หลาย

ท่านจะคิดทันทีว่าถ้าเช่นนั้นเราเปลี่ยนมาใช้สถิติไม่อิงพารามิเตอร์ดีกว่าหรือไม่เพราะไม่ต้องคำนึงถึงข้อตกลงเบื้องต้น ก่อนที่เราจะตัดสินใจว่าจะเลือกใช้สถิติอิงพารามิเตอร์ หรือสถิติไม่อิงพารามิเตอร์ ให้มาพิจารณาข้อดีและข้อด้อยของสถิติไม่อิงพารามิเตอร์ก่อน

#### **ข้อดีของสถิติไม่อิงพารามิเตอร์**

1. สามารถใช้ได้กับข้อมูลที่มีมาตราวัดตั้งแต่ระดับนามบัญญัติ
2. สามารถใช้ได้กับกลุ่มตัวอย่างที่มีขนาดเล็กและขนาดใหญ่
3. การคำนวณไม่ยุ่งยากซับซ้อน
4. ไม่จำเป็นต้องทราบการแจกแจงของประชากร และอาจมีข้อตกลงเบื้องต้นเพียงไม่กี่ข้อสำหรับสถิติไม่อิงพารามิเตอร์ บางตัว

#### **ข้อด้อยของสถิติไม่อิงพารามิเตอร์**

1. กำลังการทดสอบ (Power of a test) เมื่อข้อมูลไม่เป็นไปตามข้อตกลงเบื้องต้นสำหรับสถิติอิงพารามิเตอร์ การใช้สถิติไม่อิงพารามิเตอร์ จะให้กำลังการทดสอบสูงกว่า ในทางกลับกันถ้าข้อมูลเป็นตามข้อตกลงของการใช้สถิติอิงพารามิเตอร์ การใช้สถิติไม่อิงพารามิเตอร์ จะให้กำลังการทดสอบต่ำกว่าการใช้สถิติอิงพารามิเตอร์
2. ถ้ากลุ่มตัวอย่างที่ใช้มีขนาดใหญ่ การทดสอบโดยใช้สถิติไม่อิงพารามิเตอร์ จะมีความยุ่งยาก ต้องใช้เวลาในการคำนวณมากขึ้น ซึ่งปัญหานี้ในปัจจุบันไม่เกิดขึ้นถ้าเราใช้โปรแกรมคอมพิวเตอร์ ไม่ว่าจะเป็นโปรแกรมสำเร็จรูป หรือเขียนโปรแกรมเอง และประสิทธิภาพจะลดลงเมื่อเทียบกับการใช้สถิติอิงพารามิเตอร์

ตารางที่ 1 การเปรียบเทียบการทดสอบแบบอิงพารามิเตอร์ และแบบไม่อิงพารามิเตอร์

| การทดสอบ                                                          | การทดสอบแบบอิงพารามิเตอร์ | การทดสอบแบบไม่อิงพารามิเตอร์                                                                                                                                                                                                       |
|-------------------------------------------------------------------|---------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| การทดสอบเปรียบเทียบ<br>กรณีกลุ่มตัวอย่างเดียว                     | t – test                  | <ul style="list-style-type: none"> <li>- การทดสอบแบบรันส์</li> <li>- การทดสอบทวินาม</li> <li>- การทดสอบภาวะสารถี</li> <li>- การทดสอบคอลโมโกรอฟ-สมิรโนฟ</li> </ul> กรณีประชากรเดียว                                                 |
| การทดสอบเปรียบเทียบ<br>กรณีตัวอย่างสองกลุ่มที่เป็นอิสระกัน        | Independent<br>t – test   | <ul style="list-style-type: none"> <li>- การทดสอบไคกำลังสองสำหรับตัวอย่างสองกลุ่มที่เป็นอิสระกัน</li> <li>- การทดสอบโดยมัธยฐาน</li> <li>- การทดสอบแมนน์-วิทนี</li> <li>- การทดสอบคอลโมโกรอฟ-สมิรโนฟ</li> </ul> กรณีประชากร 2 กลุ่ม |
| การทดสอบเปรียบเทียบ<br>ตัวอย่างสองกลุ่มที่สัมพันธ์กัน             | Dependent<br>t – test     | <ul style="list-style-type: none"> <li>- การทดสอบแม็กนิตาร์</li> <li>- การทดสอบด้วยเครื่องหมาย</li> <li>- การทดสอบลำดับที่โดยเครื่องหมายของวิลค็อกซัน</li> </ul>                                                                   |
| การทดสอบเปรียบเทียบ<br>กรณีตัวอย่างมากกว่าสองกลุ่มที่เป็นอิสระกัน | One way ANOVA             | <ul style="list-style-type: none"> <li>- การทดสอบไคกำลังสองสำหรับตัวอย่างกลุ่มที่เป็นอิสระจากกัน</li> <li>- การทดสอบโดยมัธยฐาน</li> <li>- การทดสอบครุสคัล-วอลลิส</li> </ul>                                                        |
| การทดสอบเปรียบเทียบ<br>กรณีตัวอย่างมากกว่าสองกลุ่มที่สัมพันธ์กัน  | Repeated One way<br>ANOVA | <ul style="list-style-type: none"> <li>- การทดสอบของคีอกแครนคิว</li> <li>- การทดสอบของฟรีดแมน</li> </ul>                                                                                                                           |

มีข้อสังเกตว่าการทดสอบแบบอิงพารามิเตอร์จะมีตัวสถิติทดสอบเพียงตัวเดียว ตัวอย่างเช่น การทดสอบสมมติฐานเกี่ยวกับตัวอย่าง 1 กลุ่ม จะใช้ One sample t-test ( $H_0: \mu = c$ ) ส่วนการทดสอบไม่อิงพารามิเตอร์ จะมีหลายตัวที่สามารถใช้ได้ขึ้นกับสิ่งที่สนใจทดสอบ เช่นการทดสอบภาวะสสารูปดี (Goodness of fit test) ใช้ทดสอบว่า ข้อมูลที่เก็บรวบรวมมาได้มีรูปแบบสอดคล้องกับตัวแบบที่ผู้วิจัยคาดหวังหรือไม่ เช่นเราต้องการทดสอบว่าพันธุ์พืชที่ปรับปรุงใหม่มีลักษณะพันธุ์กรรมเป็นไปตามกฎของเมนเดลหรือไม่ ส่วนการทดสอบแบบรันส์เป็นการทดสอบว่าข้อมูลที่เก็บรวบรวมมามีลักษณะสุมหรือไม่ ที่เป็นเช่นนี้เพราะสถิติไม่อิงพารามิเตอร์ไม่สนใจเกี่ยวกับพารามิเตอร์และลักษณะการแจกแจง นักวิจัยที่ต้องการจะใช้สถิติไม่อิงพารามิเตอร์จึงควรศึกษาว่าสถิติไม่อิงพารามิเตอร์ใดใช้ทดสอบสมมติฐานเกี่ยวกับเรื่องใด สำหรับในบทความนี้จะแนะนำสถิติไม่อิงพารามิเตอร์ 2 ตัว คือ การทดสอบแบบรันส์ (Runs test) และการทดสอบการทดสอบคอลโมโกรอฟ-สมิรโนฟ กรณีสเปซการเดียว (One-sample Kolmogorov-Smirnov test) เนื่องจากสถิติไม่อิงพารามิเตอร์ทั้งสองนี้มีประโยชน์เพื่อเราจะได้นำไปใช้ทดสอบข้อตกลงเบื้องต้นของการทดสอบอิงพารามิเตอร์ ข้อตกลงเบื้องต้นที่สำคัญของการใช้การทดสอบ t-test คือ ข้อมูลที่นำมาศึกษาต้องมีลักษณะสุมและการแจกแจงของข้อมูลต้องมีการแจกแจงแบบปกติ

**การทดสอบแบบรันส์ (Run test)** เป็นการทดสอบว่าข้อมูลที่เก็บรวบรวมมานั้นมีลักษณะสุมหรือไม่ เงื่อนไขการทดสอบแบบวิ่งคือตัวอย่างที่เก็บมาจะต้องมาจากประชากรชุดเดียวกันและเป็นอิสระกัน การทดสอบรันส์สามารถใช้ได้กับข้อมูลตั้งแต่มาตราวัดนามบัญญัติ

**ตัวอย่างที่ 1** จากการสอบถามลูกค้าจำนวน 25 คนว่ามีความพอใจกับการให้บริการของธนาคารแห่งหนึ่งหรือไม่ (Y แทน พอใจ และ N แทน ไม่พอใจ) ได้ผลดังนี้ YYY NN YYY N Y NNN YYY NN YYY NN ต้องการทดสอบว่าข้อมูลมีลักษณะสุมหรือไม่

วิธีทำ ขั้นตอนที่ 1 ตั้งสมมติฐาน  $H_0$ : ข้อมูลมีลักษณะสุม  
 $H_1$ : ข้อมูลไม่มีลักษณะสุม

ขั้นตอนที่ 2 นับจำนวนรันส์ จะได้  $r = 10$  (จำนวนรันส์ คือ จำนวนชุดของสิ่งที่มีลักษณะเดียวกันซึ่งเรียงต่อเนื่องกัน จากข้อมูลข้างต้น YYY เป็น 1 ชุด ถัดไป NN เป็นอีก 1 ชุด ต่อไปเรื่อย ๆ จนถึง NN เป็นชุดที่ 10 )

$n_1 = 15$  จำนวนคนที่ตอบพอใจ  
 $n_2 = 10$  จำนวนคนที่ตอบไม่พอใจ

ขั้นตอนที่ 3 เปรียบเทียบค่า  $r$  กับค่าวิกฤตในตาราง เนื่องจาก ค่า  $n_1$  และค่า  $n_2$  น้อยกว่าหรือเท่ากับ 20 จะปฏิเสธสมมติฐาน  $H_0$  ถ้าค่า  $r$  มีค่าน้อยกว่าหรือเท่ากับค่าวิกฤตต่ำ (Lower critical value) หรือ ถ้าค่า  $r$  มีค่ามากกว่าหรือเท่ากับค่าวิกฤตสูง (Upper critical value) จากตารางค่าวิกฤตได้ค่าวิกฤตต่ำมีค่าเท่ากับ 7 และค่าวิกฤตสูงมีค่าเท่ากับ 18 ที่ระดับนัยสำคัญ 0.05 ดังนั้น  $r = 10$  เป็นค่าระหว่าง 7 และ 18 ( $7 < 10 < 18$ ) เราจึงยอมรับสมมติฐาน  $H_0$  และสรุปได้ว่า ข้อมูลความพึงพอใจลูกค้าข้างต้นมีลักษณะสุ่ม

ในกรณีที่ข้อมูลที่เกิดขึ้นรวบรวมมาได้มีมาตรวัดแบบช่วง หรือมาตรวัดแบบอัตราส่วน (interval scale หรือ ratio scale) การใช้การทดสอบรันส์จะต้องมีจุดอ้างอิง ซึ่งส่วนมากนิยมใช้มัธยฐาน ข้อมูลที่มีค่ามากกว่ามัธยฐานจะใช้สัญลักษณ์ + และข้อมูลที่มีค่าน้อยกว่ามัธยฐานจะใช้สัญลักษณ์ - จากนั้นนับจำนวนรันส์ของข้อมูล เพื่อเปรียบเทียบกับค่าวิกฤต และนำไปสู่ผลสรุป

สำหรับกรณีที่จำนวนข้อมูลที่มีค่า  $n_1$  หรือ ค่า  $n_2$  มากกว่า 20

$$\text{ตัวสถิติทดสอบคือ } Z = \frac{r - \left( \frac{2n_1n_2}{n_1+n_2} + 1 \right)}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1+n_2)^2(n_1+n_2-1)}}}$$

บริเวณปฏิเสธสมมติฐานหลัก คือ  $Z > z_{\alpha/2}$  หรือ  $Z < -z_{\alpha/2}$  การทดสอบรันส์เป็นการทดสอบสองด้าน

**ตัวอย่างที่ 2** คะแนนสอบนักศึกษาในกลุ่มหนึ่งจำนวน 50 คน เป็นดังนี้

|    |    |    |    |    |    |    |    |    |    |
|----|----|----|----|----|----|----|----|----|----|
| 30 | 21 | 60 | 42 | 29 | 27 | 47 | 29 | 38 | 50 |
| 32 | 40 | 56 | 41 | 61 | 28 | 37 | 42 | 57 | 20 |
| 62 | 80 | 26 | 29 | 30 | 32 | 31 | 29 | 42 | 38 |
| 42 | 23 | 23 | 54 | 29 | 32 | 41 | 49 | 43 | 45 |
| 64 | 39 | 51 | 57 | 39 | 47 | 49 | 53 | 39 | 31 |

พิจารณาว่าคะแนนนักศึกษาที่เก็บรวบรวมมามีลักษณะสุ่มหรือไม่

วิธีทำ ขั้นตอนที่ 1 ตั้งสมมติฐาน  $H_0$ : คะแนนของนักศึกษามีลักษณะสุ่ม

$H_1$ : คะแนนของนักศึกษาไม่มีลักษณะสุ่ม

ขั้นตอนที่ 2 หาค่ามัธยฐาน ได้เท่ากับ 39.5

|   |   |   |   |   |   |   |   |   |   |
|---|---|---|---|---|---|---|---|---|---|
| - | - | + | + | - | - | + | - | - | + |
| - | - | + | + | + | - | - | + | + | - |
| + | + | - | - | - | - | - | - | + | - |
| + | - | - | + | - | - | + | + | + | + |
| + | - | + | + | - | + | + | + | - | - |

จะได้จำนวนรันส์เท่ากับ 25 ( $r = 25$ )  $n_1 = 26$   $n_2 = 24$

$$Z = \frac{r - \left( \frac{2n_1n_2 + 1}{n_1 + n_2} \right)}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2(n_1 + n_2 - 1)}}}$$

$$= -0.4266$$

ค่าวิกฤตเท่ากับ -1.96 และ 1.96 ดังนั้นยอมรับสมมติฐานหลักว่าข้อมูลที่เกิดขึ้นได้มีลักษณะสุ่ม

สังเกตว่าตัวสถิติทดสอบเป็นสูตรที่ยุ่งยากแต่เมื่อขนาดตัวอย่างใหญ่ขึ้น สูตรของตัวสถิติทดสอบจะอยู่ในรูปที่ไม่ซับซ้อน

นั่นคือตัวสถิติทดสอบจะอยู่ในรูปดังต่อไปนี้

$$Z = \frac{r - E(r)}{S(r)}$$

เมื่อ  $E(r) = \frac{n+2}{2}$  และ  $S(r) = \sqrt{\frac{n(n-2)}{4(n-1)}}$  และถ้า  $n$  มากพอจะ

ประมาณ  $E(r)$  ด้วย  $\frac{n}{2}$  และ ประมาณ  $S(r)$  ด้วย  $\sqrt{\frac{n-1}{4}}$

เช่นจากตัวอย่างที่ 2  $E(r) = \frac{50+2}{2} = 26$  และ  $S(r) = \sqrt{\frac{50(50-2)}{4(50-1)}}$

$$= 3.4992$$

$$Z = \frac{25-26}{3.4992} = -0.2858$$

ซึ่งปัจจุบันปัญหาการคำนวณเหล่านี้จะไม่เกิดถ้าผู้วิจัยสามารถใช้โปรแกรมสำเร็จรูปสำหรับบทความนี้จะอ้างอิงโปรแกรมสำเร็จรูป PASW (โปรแกรม SPSS เดิม) ด้วยคำสั่ง analyze nonparametrics Legacy Dialogs Runs

NPar Tests

Runs Test

|                         | score |
|-------------------------|-------|
| Test Value <sup>a</sup> | 40    |
| Cases < Test Value      | 25    |
| Cases >= Test Value     | 25    |
| Total Cases             | 50    |
| Number of Runs          | 25    |
| Z                       | -.286 |
| Asymp. Sig. (2-tailed)  | .775  |

a. Median

รูปที่ 1 ผลลัพธ์การทดสอบรันส์ของตัวอย่างที่ 2 ด้วยโปรแกรม PASW

จากผลลัพธ์โปรแกรม PASW ในรูปที่ 1 จะได้ค่า  $Z = -0.286$  (สังเกตค่าที่ได้จะเท่ากับค่าที่ได้จากตัวสถิติทดสอบ กรณีที่ตัวอย่างมีขนาดใหญ่ และได้ค่าพี (p-value) เท่ากับ 0.775 นั้นหมายความว่าข้อมูลที่เก็บรวบรวมมาได้มีลักษณะสุ่ม

การทดสอบ คอลโมโกรอฟ-สมิร์นอฟ กรณีประชากรเดียว (Kolmogorov-Smirnov one-sample) เป็นการทดสอบว่าข้อมูลที่เก็บรวบรวมมานั้นมีรูปแบบการแจกแจงเป็นไปตามที่คาดหวังไว้หรือไม่ ข้อมูลที่จะนำมาทดสอบเป็นข้อมูลที่มีมาตราวัดตั้งแต่มาตราวัดนามบัญญัติ ด้วยเหตุนี้เราจึงสามารถใช้ การทดสอบ คอลโมโกรอฟ-สมิร์นอฟ กรณีประชากรเดียว ทดสอบว่าข้อมูลที่เก็บรวบรวมมาได้มีการแจกแจงปกติหรือไม่

ตัวอย่างที่ 3 จากตัวอย่างที่ 2 ต้องการทดสอบว่าคะแนนของนักศึกษาทั้ง 50 คนที่เก็บรวบรวมมาได้มีการแจกแจงปกติหรือไม่

วิธีทำ ขั้นตอนที่ 1 ตั้งสมมติฐาน  $H_0$ : ข้อมูลมีการแจกแจงปกติ

$H_1$ : -ข้อมูลมีการแจกแจงที่ไม่ใช่การแจกแจงปกติ

ขั้นตอนที่ 2 คำนวณหาค่าความถี่สะสมสัมพัทธ์ ( $F_0(x)$ )

จากข้อมูลที่กำหนดให้ได้ค่าเฉลี่ย ( $\bar{x}$ ) = 40.72 และค่าเบี่ยงเบนมาตรฐาน ( $S.D.$ ) = 12.96 และจัดเรียงข้อมูลจากน้อยไปมากและสร้างตารางการแจกแจงความถี่และความถี่สะสมได้ดังนี้

| คะแนน | ความถี่ | ความถี่สะสม | คะแนน | ความถี่ | ความถี่สะสม |
|-------|---------|-------------|-------|---------|-------------|
| 20    | 1       | 1           | 43    | 1       | 33          |
| 21    | 1       | 2           | 45    | 1       | 34          |
| 23    | 2       | 4           | 47    | 2       | 36          |
| 26    | 1       | 5           | 49    | 2       | 38          |
| 27    | 1       | 6           | 50    | 1       | 39          |
| 28    | 1       | 7           | 51    | 1       | 40          |
| 29    | 5       | 12          | 53    | 1       | 41          |
| 30    | 2       | 14          | 54    | 1       | 42          |
| 31    | 2       | 16          | 56    | 1       | 43          |
| 32    | 3       | 19          | 57    | 2       | 45          |
| 37    | 1       | 20          | 60    | 1       | 46          |
| 38    | 2       | 22          | 61    | 1       | 47          |
| 39    | 3       | 25          | 62    | 1       | 48          |
| 40    | 1       | 26          | 64    | 1       | 49          |
| 41    | 2       | 28          | 80    | 1       | 50          |
| 42    | 4       | 32          |       |         |             |

เปลี่ยนคะแนนให้เป็นค่ามาตรฐาน โดย  $Z = \frac{X - \bar{X}}{S.D.}$  และคำนวณค่าความถี่สะสมสัมพัทธ์



| คะแนน | ค่ามาตรฐาน | $F_O(x)$ | คะแนน | ค่ามาตรฐาน | $F_O(x)$ |
|-------|------------|----------|-------|------------|----------|
| 20    | -1.5988    | 0.02     | 43    | 0.1759     | 0.66     |
| 21    | -1.5216    | 0.04     | 45    | 0.3302     | 0.68     |
| 23    | -1.3673    | 0.08     | 47    | 0.4846     | 0.72     |
| 26    | -1.1358    | 0.1      | 49    | 0.6389     | 0.76     |
| 27    | -1.0586    | 0.12     | 50    | 0.7160     | 0.78     |
| 28    | -0.9815    | 0.14     | 51    | 0.7932     | 0.8      |
| 29    | -0.9043    | 0.24     | 53    | 0.9475     | 0.82     |
| 30    | -0.8272    | 0.28     | 54    | 1.0247     | 0.84     |
| 31    | -0.7500    | 0.32     | 56    | 1.1790     | 0.86     |
| 32    | -0.6728    | 0.38     | 57    | 1.2562     | 0.9      |
| 37    | -0.2870    | 0.4      | 60    | 1.4877     | 0.92     |
| 38    | -0.2099    | 0.44     | 61    | 1.5648     | 0.94     |
| 39    | -0.1327    | 0.5      | 62    | 1.6420     | 0.96     |
| 40    | -0.0556    | 0.52     | 64    | 1.7963     | 0.98     |
| 41    | 0.0216     | 0.56     | 80    | 3.0309     | 1        |
| 42    | 0.0988     | 0.64     |       |            |          |

ขั้นตอนที่ 3 หาค่า  $P(Z < z) = F_E(x)$  จากตารางการแจกแจงปกติมาตรฐาน โดยจากสมมติฐานหลักที่คาดว่าข้อมูลที่เก็บรวบรวมมาได้มีการแจกแจงปกติ ดังนั้นควรจะได้ว่า  $P(Z < z)$  และ  $F_O(x)$  ควรมีค่าใกล้เคียงกัน หรือต่างกันไม่มาก

| $z$     | $P(Z < z)$<br>$= F_E(x)$ | $F_O(x)$ | $z$    | $P(Z < z)$<br>$= F_E(x)$ | $F_O(x)$ |
|---------|--------------------------|----------|--------|--------------------------|----------|
| -1.5988 | 0.1111                   | 0.02     | 0.1759 | 0.6072                   | 0.66     |
| -1.5216 | 0.1254                   | 0.04     | 0.3302 | 0.6222                   | 0.68     |
| -1.3673 | 0.1567                   | 0.08     | 0.4846 | 0.6452                   | 0.72     |
| -1.1358 | 0.2093                   | 0.1      | 0.6389 | 0.6747                   | 0.76     |
| -1.0586 | 0.2278                   | 0.12     | 0.7160 | 0.6913                   | 0.78     |
| -0.9815 | 0.2465                   | 0.14     | 0.7932 | 0.7087                   | 0.8      |
| -0.9043 | 0.2651                   | 0.24     | 0.9475 | 0.7453                   | 0.82     |
| -0.8272 | 0.2834                   | 0.28     | 1.0247 | 0.7640                   | 0.84     |
| -0.7500 | 0.3011                   | 0.32     | 1.1790 | 0.8009                   | 0.86     |
| -0.6728 | 0.3181                   | 0.38     | 1.2562 | 0.8188                   | 0.9      |
| -0.2870 | 0.3828                   | 0.4      | 1.4877 | 0.8681                   | 0.92     |
| -0.2099 | 0.3903                   | 0.44     | 1.5648 | 0.8827                   | 0.94     |
| -0.1327 | 0.3954                   | 0.5      | 1.6420 | 0.8964                   | 0.96     |
| -0.0556 | 0.3983                   | 0.52     | 1.7963 | 0.9205                   | 0.98     |
| 0.0216  | 0.6012                   | 0.56     | 3.0309 | 0.9960                   | 1        |
| 0.0988  | 0.6030                   | 0.64     |        |                          |          |

ขั้นตอนที่ 4 คำนวณหาค่าความแตกต่างระหว่าง  $F_O(x)$  กับ  $F_E(x)(D_{cal})$  และทำการเปรียบเทียบค่าความแตกต่างที่มากที่สุดที่คำนวณได้ ( $\max(D_{cal})$ ) กับ ค่า  $D_{table}$  จากตาราง Kolmogorov-Smirnov one-sample ถ้า  $\max(D_{cal}) \geq D_{table}$  ปฏิเสธ  $H_0: F_0 = F_E$  หมายความว่าข้อมูลมีการแจกแจงที่ไม่ใช่การแจกแจงปกติ และ ถ้า  $\max(D_{cal}) < D_{table}$  ยอมรับ  $H_0: F_0 = F_E$  หมายความว่าข้อมูลมีการแจกแจงปกติ

| $F_O(x)$ | $P(Z < z)$<br>$= F_E(x)$ | $D_{cal}$ | $F_O(x)$ | $P(Z < z)$<br>$= F_E(x)$ | $D_{cal}$ |
|----------|--------------------------|-----------|----------|--------------------------|-----------|
| 0.02     | 0.1111                   | 0.0911    | 0.66     | 0.6072                   | 0.0528    |
| 0.04     | 0.1254                   | 0.0854    | 0.68     | 0.6222                   | 0.0578    |
| 0.08     | 0.1567                   | 0.0767    | 0.72     | 0.6452                   | 0.0748    |
| 0.1      | 0.2093                   | 0.1093    | 0.76     | 0.6747                   | 0.0853    |
| 0.12     | 0.2278                   | 0.1078    | 0.78     | 0.6913                   | 0.0887    |
| 0.14     | 0.2465                   | 0.1065    | 0.8      | 0.7087                   | 0.0913    |
| 0.24     | 0.2651                   | 0.0251    | 0.82     | 0.7453                   | 0.0747    |
| 0.28     | 0.2834                   | 0.0034    | 0.84     | 0.7640                   | 0.0760    |
| 0.32     | 0.3011                   | 0.0189    | 0.86     | 0.8009                   | 0.0591    |
| 0.38     | 0.3181                   | 0.0619    | 0.9      | 0.8188                   | 0.0812    |
| 0.4      | 0.3828                   | 0.0172    | 0.92     | 0.8681                   | 0.0519    |
| 0.44     | 0.3903                   | 0.0497    | 0.94     | 0.8827                   | 0.0573    |
| 0.5      | 0.3954                   | 0.1046    | 0.96     | 0.8964                   | 0.0636    |
| 0.52     | 0.3983                   | 0.1217    | 0.98     | 0.9205                   | 0.0595    |
| 0.56     | 0.6012                   | 0.0412    | 1        | 0.9960                   | 0.0040    |
| 0.64     | 0.6030                   | 0.0370    |          |                          |           |

จากตารางข้างต้น  $\max(D_{cal}) = 0.1217$  ที่ระดับนัยสำคัญ 0.05 ได้  
ค่า  $D_{table} = \frac{1.36}{\sqrt{n}} = \frac{1.36}{\sqrt{50}} = 0.1923$  สรุปยอมรับสมมติฐาน  $H_0: F_0 = F_E$   
นั่นแสดงว่าคะแนนสอบของนักศึกษา 50 คนมีการแจกแจงปกติ

การทดสอบคอลโมโกรอฟ-สมิธโนฟ กรณีประชากรเดียว มีขั้นตอนการคำนวณที่ยุ่งยาก  
มาก นักวิจัยอาจใช้โปรแกรมสำเร็จรูปช่วยซึ่งโปรแกรมสำเร็จรูปอาจมีหลายโปรแกรมเช่นเดียวกับการ  
ทดสอบแบบรันส์ แต่สำหรับบทความนี้จะอ้างอิงโปรแกรมสำเร็จรูป PASW (โปรแกรม SPSS  
เดิม) ด้วยคำสั่ง คำสั่ง analyze nonparametrics Legacy Dialogs 1-sample K-S

NPar Tests

One-Sample Kolmogorov-Smirnov Test

|                                  |                | score  |
|----------------------------------|----------------|--------|
| N                                |                | 50     |
| Normal Parameters <sup>a,b</sup> | Mean           | 40.72  |
|                                  | Std. Deviation | 12.963 |
| Most Extreme Differences         | Absolute       | .129   |
|                                  | Positive       | .129   |
|                                  | Negative       | -.055  |
| Kolmogorov-Smirnov Z             |                | .915   |
| Asymp. Sig. (2-tailed)           |                | .372   |

a. Test distribution is Normal.

b. Calculated from data.

รูปที่ 2 ผลลัพธ์การทดสอบคอลโมโกรอฟ-สมิรโนฟ กรณีประชากรเดียว ของตัวอย่างที่ 3 ด้วยโปรแกรม PASW

จากผลลัพธ์โปรแกรม PASW ในรูปที่ 2 ได้ค่าพี (p-value) เท่ากับ 0.372 ดังนั้นยอมรับสมมติฐาน  $H_0: F_0 = F_E$  นั้นแสดงว่าคะแนนของนักศึกษาที่เก็บรวบรวมได้มีการแจกแจงปกติ

**บทสรุป** สถิติไม่อิงพารามิเตอร์เป็นสถิติวิเคราะห์ที่น่าสนใจนำมาใช้ในงานวิจัยเพราะสถิติไม่อิงพารามิเตอร์ใช้ได้กับข้อมูลทุกลักษณะการแจกแจง มาตราวัดข้อมูลเริ่มต้นตั้งแต่มาตราวัดนามบัญญัติ ขนาดตัวอย่างจะเล็กหรือใหญ่ก็ได้ และมีสถิติไม่อิงพารามิเตอร์ให้เลือกใช้หลายตัวเพื่อให้สอดคล้องกับวัตถุประสงค์ของงานวิจัย แม้ว่าสูตรที่ใช้คำนวณอาจมีความยุ่งยาก แต่ปัจจุบันผู้วิจัยสามารถใช้โปรแกรมสำเร็จรูปทางสถิติไม่ว่าจะเป็นโปรแกรม PASW (โปรแกรม SPSS เดิม) หรือโปรแกรมไม่มีลิขสิทธิ์ เช่นโปรแกรม R แต่อย่างไรก็ดีผู้วิจัยจะต้องคำนึงถึงประสิทธิภาพของสถิติทดสอบ ถ้าในกรณีที่ข้อมูลมีลักษณะผ่านข้อตกลงเบื้องต้นของสถิติอิงพารามิเตอร์ ผู้วิจัยจะต้องเลือกใช้สถิติอิงพารามิเตอร์ทดสอบ เพราะสถิติอิงพารามิเตอร์มีประสิทธิภาพสูงกว่าสถิติไม่อิงพารามิเตอร์เว้นเสียแต่ว่าข้อมูลที่จะนำมาวิเคราะห์ไม่ผ่านข้อตกลงเบื้องต้น หรือตัวอย่างมีขนาดเล็กแม้ว่าข้อมูลจะผ่านข้อตกลงเบื้องต้น ผู้วิจัยควรเลือกสถิติไม่อิงพารามิเตอร์เป็นสถิติทดสอบ

เอกสารอ้างอิง

นิภา ศรีไพโรจน์. (2533). สถิตินอนพาราเมตริก. กรุงเทพฯ : โอเดียนสโตร์.

ศิริลักษณ์ สุวรรณวงศ์. (2560). สถิติไม่ใช่พารามิเตอร์ : เอกสารคำสอนวิชา SCMA 384.

นครปฐม : คณะวิทยาศาสตร์ มหาวิทยาลัยมหิดล.

Siegel Sidney. (1981). *Nonparametric Statistics for the Behavioral Sciences*.

New York : McGraw-Hill.

ผู้เขียน

ศิริลักษณ์ สุวรรณวงศ์

คณะวิทยาศาสตร์ มหาวิทยาลัยมหิดล

E-mail: sirilak.suw@mahidol.ac.th