



ใบรับรองวิทยานิพนธ์

บัณฑิตวิทยาลัย สถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ

เรื่อง การเปรียบเทียบอำนาจการทดสอบภาวะสารูปติของสถิติบางตัว

สำหรับการแจกแจงลอการิธึมและไวบูลล์

โดย นางสาวศรีอัมพร เร่บ้านเกาะ

ได้รับอนุมัติให้นับเป็นส่วนหนึ่งของการศึกษาตามหลักสูตร

วิทยาศาสตรมหาบัณฑิต สาขาวิชาสถิติประยุกต์

คณบดีบัณฑิตวิทยาลัย

(อาจารย์ ดร.มงคล หวังสถิตย์วงศ์)

23 มีนาคม 2550

คณะกรรมการสอบวิทยานิพนธ์

ประธานกรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.วินัย โพธิ์สุวรรณ)

กรรมการ

(Dr. Gareth Clayton)

กรรมการ

(รองศาสตราจารย์อดิศักดิ์ พงษ์พูลผลศักดิ์)

กรรมการ

(รองศาสตราจารย์สอาด นิวิศพงศ์)

การพัฒนาและหาประสิทธิภาพบทเรียนคอมพิวเตอร์ช่วยสอน
เรื่องการเขียนแผนการจัดการเรียนรู้

นางสาวชนาด เล็บครุฑ

วิทยานิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร
ครุศาสตรบัณฑิต สาขาวิชาเทคโนโลยีเทคนิคศึกษา ภาควิชาครุศาสตร์เทคโนโลยี
บัณฑิตวิทยาลัย สถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ
ปีการศึกษา 2549
ลิขสิทธิ์ของสถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ

ชื่อ : นางสาวศรีอัมพร เร่บ้านเกาะ
ชื่อวิทยานิพนธ์ : การเปรียบเทียบอำนาจการทดสอบภาวะสารถีของสถิติบางตัว
สำหรับการแจกแจงลอการิทึมและไวบูลล์
สาขาวิชา : สถิติประยุกต์
สถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ
ที่ปรึกษาวิทยานิพนธ์ : ผู้ช่วยศาสตราจารย์ ดร.วินัย โพธิ์สุวรรณ
อาจารย์ ดร.กาเรธ เคลตัน
ปีการศึกษา : 2549

บทคัดย่อ

การวิจัยครั้งนี้มีวัตถุประสงค์เพื่อเปรียบเทียบตัวสถิติทดสอบที่ใช้ในการจำแนกการแจกแจงลอการิทึม (Log-normal Distribution) และการแจกแจงแบบไวบูลล์ (Weibull Distribution) ซึ่งสถิติทดสอบที่ใช้ในการวิจัย คือ Anderson Darling Test Statistic (AD), Kuiper Test Statistic(K) , Ratio of Maximum Likelihoods Test Statistic (RML) และ Kolmogorov -Smirnov Test Statistic (KS) ในการวิจัยครั้งนี้กำหนดลักษณะการแจกแจงของประชากรเป็นการแจกแจงลอการิทึม 2 พารามิเตอร์และ การแจกแจงไวบูลล์ 2 พารามิเตอร์ ด้วยขนาดตัวอย่าง 20, 30, 40 และ 50 ทำการทดลอง ซ้ำๆ กัน 1,000 ครั้ง ในแต่ละสถานการณ์ โดยพิจารณาจาก ความสามารถในการควบคุมค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 แล้วทำการเปรียบเทียบประสิทธิภาพของวิธีการทดสอบด้วยค่าอำนาจของ การทดสอบของสถิติทดสอบ 4 วิธี ภายใต้ระดับนัยสำคัญ (α) 0.01, 0.05 และ 0.10 ซึ่งผลการวิจัยสรุปได้ดังนี้ การพิจารณาจากความสามารถในการควบคุมค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยใช้เกณฑ์ของ Cochran ผลปรากฏว่า ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ดีที่สุด และโดยใช้เกณฑ์ของ Bradley ผลปรากฏว่า ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ดีที่สุด ผลการเปรียบเทียบค่าอำนาจของการทดสอบของตัวสถิติทดสอบดังกล่าวผลปรากฏว่า ค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าต่ำสุด และใกล้เคียงกัน

(วิทยานิพนธ์มีจำนวนทั้งสิ้น 133 หน้า)

คำสำคัญ : การทดสอบภาวะสารถี, การแจกแจงลอการิทึม, การแจกแจงแบบไวบูลล์

อาจารย์ที่ปรึกษาวิทยานิพนธ์

Name : Miss Sriamporn Raibankoh
Thesis Title : Comparison of Power of Some Goodness of Fit Tests for Testing
Lognormal and Weibull Distribution
Major Field : Applied Statistics
King Mongkut's Institute of technology North Bangkok
Thesis Advisors : Assistant Professor Dr.Winai Bodhisuwan
Dr.Gareth Clayton
Academic Year : 2006

Abstract

The objective of this research was to compare on the power of some goodness of fit test statistics for Log-normal distribution and Weibull distribution. The test statistics are the Anderson Darling Test Statistic (AD), the Kuiper Test Statistic (K), the Ration of Maximum Likelihoods Test Statistic (RML) and the Kolmogorov-Smirnov Test Statistic (KS). The data for this research were generated from the two parameter Lognormal Distribution and two parameter Weibull Distribution and the experiment was repeated 1000 times. All parameter combinations were done with 20, 30, 40 and 50. Their efficiencies were compared by controlling the probability of type I error and power of some goodness of fit tests. The results of this study show that the Ratio of Maximum Likelihoods Test Statistic (RML) best controls the probability of type I error best fit for the largest number data sets using the Cochran criterion and the Kolmogorov -Smirnov Test Statistic (KS) best controls the probability of type I error for the largest number data sets using the Bradley criterion. The conclusion of research comparing the power of some goodness of fit test is that Ratio of Maximum Likelihoods Test Statistic (RML) has highest power in all situations

(Total 133 pages)

Keywords : Goodness of fit test, Log-normal distribution, Weibull distribution

Advisor

กิตติกรรมประกาศ

วิทยานิพนธ์ฉบับนี้สำเร็จลุล่วงไปได้ความช่วยเหลืออย่างดียิ่งของ ผู้ช่วยศาสตราจารย์ ดร.วินัย โพธิ์สุวรรณ อาจารย์ที่ปรึกษาวิทยานิพนธ์หลัก และดร.กาเรธ เคลตัน อาจารย์ที่ปรึกษาวิทยานิพนธ์ร่วม ที่ได้ให้คำปรึกษาและข้อเสนอแนะ ตลอดการศึกษาและดำเนินการวิจัย ผู้วิจัยขอกราบขอบพระคุณเป็นอย่างสูงไว้ ณ โอกาสนี้

ขอขอบคุณรองศาสตราจารย์อดิศักดิ์ พงษ์พูลผลศักดิ์ และรองศาสตราจารย์สอาด นิเวศพงศ์ ที่สละเวลามาเป็นกรรมการในการสอบวิทยานิพนธ์ฉบับนี้

ขอขอบคุณอาจารย์ในสาขาคณิตศาสตร์ทุกท่าน ของวิทยาลัยเทคโนโลยีอุตสาหกรรม สถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ ที่ให้ความช่วยเหลือในทุกๆ ด้านและเป็นกำลังใจในการทำงานวิจัยมาโดยตลอด

ท้ายนี้ผู้วิจัยใคร่ขอกราบขอบพระคุณบิดา มารดาที่ส่งเสริมสนับสนุนการศึกษาในทุกๆ ด้าน รวมทั้งให้ความรักความเข้าใจและห่วงใยผู้วิจัยเสมอมา ขอขอบคุณเพื่อนๆ พี่ๆ น้องๆ ที่ให้ความช่วยเหลือและกำลังใจตลอดงานวิจัยครั้งนี้

ศรีอัมพร เร่บ้านเกาะ

สารบัญ

	หน้า
บทคัดย่อภาษาไทย	ข
บทคัดย่อภาษาอังกฤษ	ค
กิตติกรรมประกาศ	ง
สารบัญตาราง	ช
สารบัญภาพ	ญ
คำอธิบายสัญลักษณ์และคำย่อ	ณ
บทที่ 1 บทนำ	1
1.1 ที่มาและความสำคัญของปัญหา	1
1.2 วัตถุประสงค์ในการวิจัย	4
1.3 สมมติฐานของการวิจัย	4
1.4 ขอบเขตของการวิจัย	4
1.5 เกณฑ์ที่ใช้ในการศึกษาประสิทธิภาพของตัวแบบ	6
1.6 ประโยชน์ของการวิจัยที่คาดว่าจะได้รับ	6
บทที่ 2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง	7
2.1 การแจกแจงของข้อมูลที่ใช้ในการวิจัย	7
2.2 การเปรียบเทียบการแจกแจงและอัตราการชำรุดของการแจกแจงลอกนอร์มัลและการแจกแจงไวบูลล์	8
2.3 การทดสอบภาวะสภาวะปกติ	11
2.4 ตัวสถิติทดสอบที่ใช้ในการวิจัย	11
2.5 ผลงานวิจัยที่เกี่ยวข้อง	21
บทที่ 3 การดำเนินการวิจัย	25
3.1 แผนการทดลอง	25
3.2 ขั้นตอนในการวิจัย	26
3.3 ขั้นตอนการทำงานของโปรแกรม	30
บทที่ 4 ผลการวิจัย	35
4.1 การเปรียบเทียบความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1	37
4.2 การเปรียบเทียบอำนาจการทดสอบ	68
บทที่ 5 สรุปผลการวิจัย	99
5.1 สรุปผลการวิจัย	99
5.2 ปัญหาที่เกิดขึ้นในการวิจัย	101

สารบัญ(ต่อ)

	หน้า
5.3 ข้อเสนอแนะ	101
บรรณานุกรม	103
ภาคผนวก ก โปรแกรมที่ใช้ในการสร้างการแจกแจงที่ใช้ในการวิจัย	105
ภาคผนวก ข โปรแกรมในการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่1	109
ภาคผนวก ค โปรแกรมในการหาค่าอำนาจการทดสอบ	121
ประวัติผู้วิจัย	133

สารบัญตาราง

ตารางที่	หน้า
2-1 การเปรียบเทียบการแจกแบบลอการิธึมและการแจกแจงแบบไวบูลล์	9
2-2 ค่า A_c สำหรับ H_0 : การแจกแจงลอการิธึมและ H_1 : การแจกแจงไวบูลล์	12
2-3 ค่า A_c สำหรับ H_0 : การแจกแจงไวบูลล์และ H_1 : การแจกแจงลอการิธึม	13
2-4 ค่า D_c สำหรับ H_0 : การแจกแจงลอการิธึมและ H_1 : การแจกแจงไวบูลล์	15
2-5 ค่า D_c สำหรับ H_0 : การแจกแจงไวบูลล์และ H_1 : การแจกแจงลอการิธึม	15
2-6 ค่า V_c สำหรับ H_0 : การแจกแจงลอการิธึมและ H_1 : การแจกแจงไวบูลล์	17
2-7 ค่า V_c สำหรับ H_0 : การแจกแจงไวบูลล์และ H_1 : การแจกแจงลอการิธึม	18
2-8 ค่า $(RML)_c^{1/n}$ สำหรับ H_0 : การแจกแจงลอการิธึมและ H_1 : การแจกแจงไวบูลล์	19
2-9 ค่า $(RML)_c^{1/n}$ สำหรับ H_0 : การแจกแจงไวบูลล์และ H_1 : การแจกแจงลอการิธึม	20
4-1 แสดงความคลาดเคลื่อนในการทดสอบสมมติฐานทางสถิติ	36
4-2 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$	38
4-3 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$	41
4-4 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$	44
4-5 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$	47
4-6 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$	50
4-7 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 0.5$	53

สารบัญตาราง(ต่อ)

ตารางที่	หน้า
4-8 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$	56
4-9 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$	59
4-10 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 2$	62
4-11 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 3$	65
4-12 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$	68
4-13 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$	71
4-14 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$	74
4-15 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$	77
4-16 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$	80
4-17 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 0.5$	83
4-18 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$	86
4-19 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$	89

สารบัญตาราง(ต่อ)

ตารางที่	หน้า
4-20 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 2$	92
4-21 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 3$	95

สารบัญภาพ

ภาพที่	หน้า
1-1 รูปโค้งการแจกแจงลอกนอร์มัล และการแจกแจงไวบูลล์	2
1-2 แสดงอัตราการชำรุดของการแจกแจงลอกนอร์มัลและไวบูลล์	2
2-1 กราฟของการแจกแจงไวบูลล์ที่ $\beta=1$ และ $\alpha=0.5, 1, 1.5, 2$ และ 3	7
2-2 กราฟของการแจกแจงลอกนอร์มัลที่ $\mu=1$ และ $\sigma=0.5, 1, 1.5, 2$ และ 3	8
2-3 รูปโค้งการแจกแจงลอกนอร์มัลที่ $\mu=1$ และ $\sigma=1$ และการแจกแจงไวบูลล์ที่ $\beta=4.104$ และ $\alpha=0.889$	10
2-4 แสดงอัตราการชำรุดของการแจกแจงลอกนอร์มัล $\mu=1$ และ $\sigma=1$ และการแจกแจงไวบูลล์ที่ $\beta=4.104$ และ $\alpha=0.889$	10
2-5 แสดงฟังก์ชันการแจกแจงความถี่สัมพัทธ์สะสมของตัวอย่าง หรือความถี่สะสมที่สังเกตได้ในรูปของสัดส่วน: $S(x_i)$, ฟังก์ชันการแจกแจงความถี่สะสมภายใต้ $H_0 : F_0[x_i]$ และค่าสถิติ KS	14
3-1 แผนผังแสดงขั้นตอนการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1	31
3-2 แผนผังแสดงขั้นตอนการหาค่าอำนาจการทดสอบ	33
4-1 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1=0.1$	39
4-2 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1=0.05$	40
4-3 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1=0.01$	40
4-4 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1$ ที่ $\alpha_1=0.1$	42
4-5 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1$ ที่ $\alpha_1=0.05$	43

สารบัญภาพ (ต่อ)

ภาพที่	หน้า
4-6 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$ ที่ $\alpha_1 = 0.01$	43
4-7 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$ ที่ $\alpha_1 = 0.1$	45
4-8 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$ ที่ $\alpha_1 = 0.05$	46
4-9 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$ ที่ $\alpha_1 = 0.01$	46
4-10 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$ ที่ $\alpha_1 = 0.1$	48
4-11 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$ ที่ $\alpha_1 = 0.05$	49
4-12 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$ ที่ $\alpha_1 = 0.01$	49
4-13 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$ ที่ $\alpha_1 = 0.1$	51
4-14 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$ ที่ $\alpha_1 = 0.05$	52

สารบัญภาพ (ต่อ)

ภาพที่	หน้า
4-15 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$ ที่ $\alpha_1 = 0.01$	52
4-16 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.1$	54
4-17 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.05$	55
4-18 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.01$	55
4-19 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.1$	57
4-20 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.05$	58
4-21 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.01$	58
4-22 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.1$	60
4-23 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.05$	61

สารบัญภาพ (ต่อ)

ภาพที่	หน้า
4-24 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=1.5$ ที่ $\alpha_1 = 0.01$	61
4-25 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=2$ ที่ $\alpha_1 = 0.1$	63
4-26 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=2$ ที่ $\alpha_1 = 0.05$	64
4-27 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=2$ ที่ $\alpha_1 = 0.01$	64
4-28 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=3$ ที่ $\alpha_1 = 0.1$	66
4-29 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=3$ ที่ $\alpha_1 = 0.05$	67
4-30 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=3$ ที่ $\alpha_1 = 0.01$	67
4-31 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1 = 0.1$	69
4-32 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1 = 0.05$	69
4-33 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1 = 0.01$	70
4-34 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\sigma=1$ ที่ $\alpha_1 = 0.1$	72

สารบัญภาพ (ต่อ)

[illegible]

สารบัญภาพ (ต่อ)

ภาพที่	หน้า
4-49 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.1$	87
4-50 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.05$	87
4-51 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.01$	88
4-52 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.1$	90
4-53 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.05$	90
4-54 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.01$	91
4-55 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.1$	93
4-56 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.05$	93
4-57 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.01$	94
4-58 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.1$	96
4-59 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.05$	96
4-60 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.01$	97

คำอธิบายสัญลักษณ์และคำย่อ

α	หมายถึง shape parameter ของการแจกแจงไวบูลล์
$\hat{\alpha}$	หมายถึง ค่าที่ได้จากการประมาณพารามิเตอร์ α ของการแจกแจงไวบูลล์
α_1	หมายถึง ระดับนัยสำคัญ หรือค่าความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้
β	หมายถึง scale parameter ของการแจกแจงไวบูลล์
$\hat{\beta}$	หมายถึง ค่าที่ได้จากการประมาณพารามิเตอร์ β ของการแจกแจงไวบูลล์
β_1	หมายถึง ค่าความคลาดเคลื่อนประเภทที่ 2
σ	หมายถึง shape parameter ของการแจกแจงล็อกนอร์มัล
$\hat{\sigma}$	หมายถึง ค่าที่ได้จากการประมาณพารามิเตอร์ σ ของการแจกแจงล็อกนอร์มัล
μ	หมายถึง scale parameter ของการแจกแจงล็อกนอร์มัล
$\hat{\mu}$	หมายถึง ค่าที่ได้จากการประมาณพารามิเตอร์ μ ของการแจกแจงไวบูลล์
τ	หมายถึง ค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ที่เกิดจากการทดลอง ตัวสถิติทดสอบมีค่าเท่ากับจำนวนครั้งของการปฏิเสธสมมติฐานว่าง เมื่อสมมติฐานว่างเป็นจริงหารด้วยจำนวนครั้งในการทดลอง
-	หมายถึง กรณีศึกษาที่ไม่สามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้
A_c	หมายถึง ค่าวิกฤตของ สถิติทดสอบ Anderson Darling
AD	หมายถึง สถิติทดสอบ Anderson Darling
D_c	หมายถึง ค่าวิกฤตของ สถิติทดสอบ Kolmogorov -Smirnov
$F_0(x)$	หมายถึง เป็นฟังก์ชันการแจกแจงสะสมตามการแจกแจงที่คาดหวัง หรือ เป็นฟังก์ชันการแจกแจงความถี่สะสมภายใต้ H_0
H_0	หมายถึง สมมติฐานว่าง (the null hypothesis)
H_1	หมายถึง สมมติฐานแย้ง (the alternative hypothesis)
K	หมายถึง สถิติทดสอบ Kuiper
KS	หมายถึง สถิติทดสอบ Kolmogorov -Smirnov
n	หมายถึง ขนาดตัวอย่าง
RML	หมายถึง สถิติทดสอบ Ratio of Maximum Likelihoods
$(RML)_c^{\frac{1}{n}}$	หมายถึง ค่าวิกฤตของสถิติทดสอบ Ratio of Maximum Likelihoods
$S(x)$	หมายถึง เป็นฟังก์ชันการแจกแจงของข้อมูลที่สังเกตได้ หรือความถี่สะสมที่สังเกตได้ ในรูปของสัดส่วน (empirical distribution function)
V_c	หมายถึง คือค่าวิกฤตของ สถิติทดสอบ Kuiper

บทที่ 1

บทนำ

1.1 ที่มาและความสำคัญของปัญหา

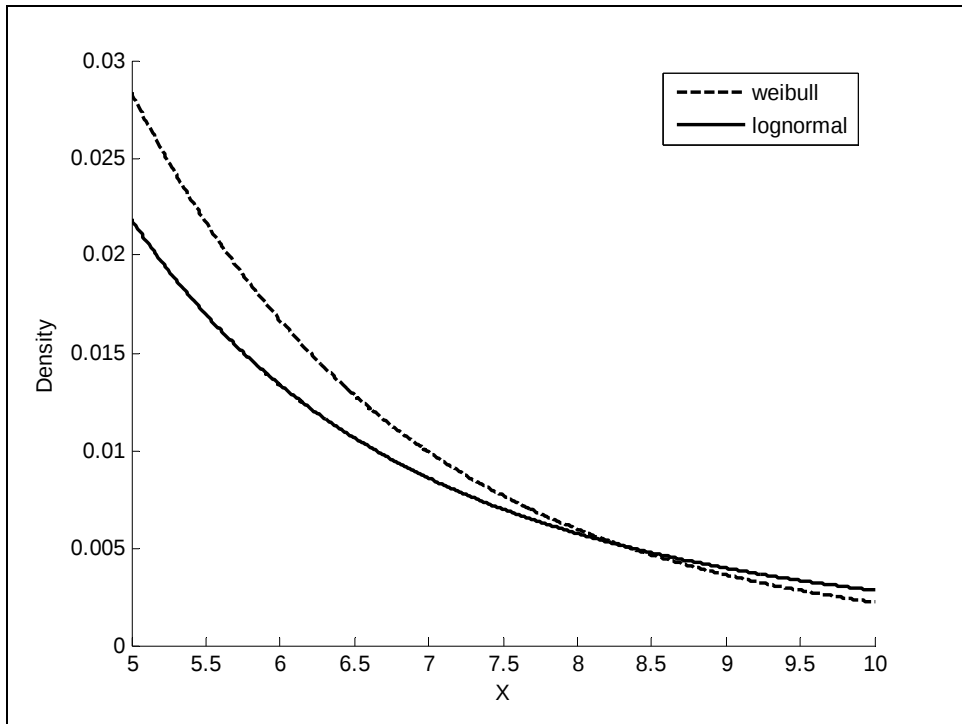
การวิเคราะห์ข้อมูลจำเป็นต้องอาศัยวิธีการทางด้านสถิติเพื่อหาผลสรุปของสมมุติฐานที่ตั้งขึ้น ข้อตกลงเบื้องต้นของตัวสถิติที่ใช้ในการวิเคราะห์ข้อมูลแบ่งเป็น 2 ประเภท คือ สถิติที่ไม่ใช้พารามิเตอร์ (Nonparametric Statistics) จะใช้วิธีการนี้เมื่อไม่สามารถระบุได้ว่าข้อมูลที่นำมาวิเคราะห์มีการแจกแจงเป็นแบบใด และสถิติที่ใช้พารามิเตอร์ (Parametric Statistics) จะใช้เพื่อวิเคราะห์ข้อมูล เมื่อตัวอย่างที่สุ่มมาจากประชากรทราบว่าการแจกแจงอย่างใดอย่างหนึ่งจะประมาณค่าของพารามิเตอร์ที่ไม่ทราบจากค่าสถิติ (Statistics)

ความน่าเชื่อถือในผลิตภัณฑ์ (Product Reliability) เป็นวิธีการสถิติที่ใช้พารามิเตอร์ในการวิเคราะห์ข้อมูล ความน่าเชื่อถือในผลิตภัณฑ์เป็นตัวบ่งชี้ที่สำคัญตัวหนึ่งของคุณภาพของผลิตภัณฑ์ ในขณะที่ผลิตภัณฑ์ก็มีรูปแบบและการปรับปรุงขึ้นทุกๆ วันและกรรมวิธีการผลิตก็ต้องละเอียดและแม่นยำขึ้น ฉะนั้นการวัดและหาความน่าเชื่อถือในผลิตภัณฑ์ก็มีความสำคัญเป็นอันมาก ความน่าเชื่อถือในผลิตภัณฑ์นี้ก็สามารถจะชี้หรือบอกคุณภาพของผลิตภัณฑ์ให้ได้ทราบอย่างชัดเจน

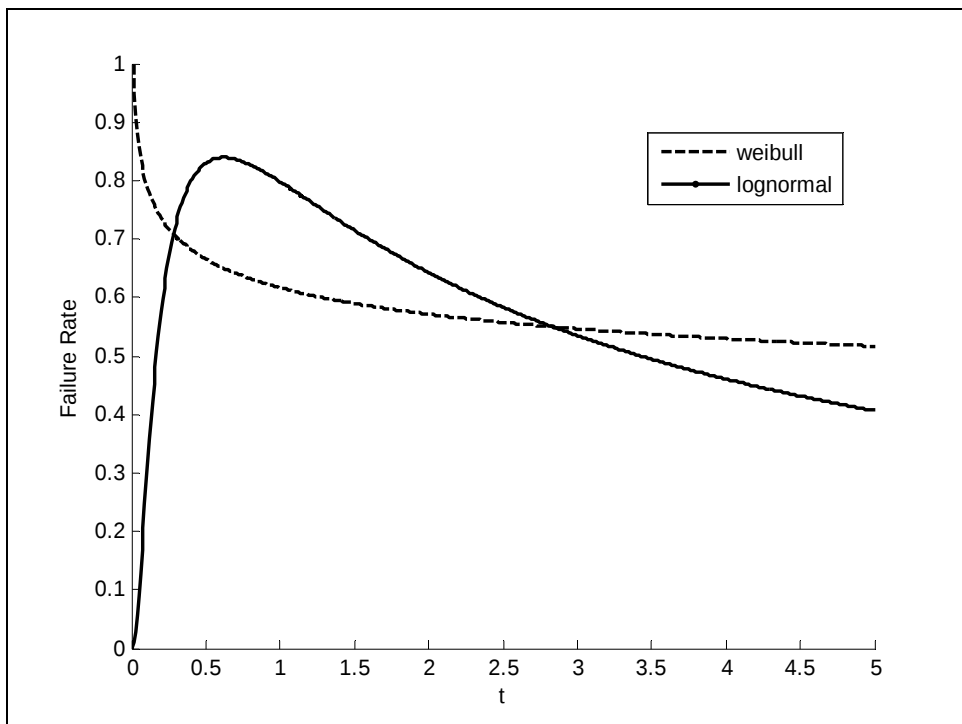
เนื่องจากความน่าเชื่อถือในผลิตภัณฑ์เป็นสถิติที่ใช้พารามิเตอร์ จึงจำเป็นต้องทราบถึงการแจกแจงของข้อมูลว่ามีการแจกแจงรูปแบบใดเพื่อวิเคราะห์หาค่าความน่าเชื่อถือในผลิตภัณฑ์ได้อย่างถูกต้องและแม่นยำ

การแจกแจงลอการิธึม (Log-normal Distribution) และการแจกแจงไวบูลล์ (Weibull Distribution) มักใช้ในการพิจารณาถึงสถานการณ์ของโค้งการแจกแจงสำหรับตัวแปรสุ่มที่เป็นบวกเสมอโดยทั้งการแจกแจงลอการิธึมและการแจกแจงไวบูลล์ มีการประยุกต์ใช้มากทางด้านวิศวกรรม

Larsen (ค.ศ.1970) เสนอตัวอย่างซึ่งเราจะเห็นได้จากการตีความค่าเฉลี่ยของข้อมูลจำนวนมากที่นำมาเปรียบเทียบ โดยเราสังเกตและพิจารณาจากการแจกแจงของมลพิษ ซึ่งมีความใกล้เคียงกับฟังก์ชันลอการิธึม ข้อมูลการแผ่รังสีมีค่าใกล้เคียงกับการแจกแจงลอการิธึมและการแจกแจงไวบูลล์ โดยมีจุดประสงค์เพื่อทำการเลือกหนึ่งในการแจกแจง โดยการทดสอบจากตัวทดสอบทางสถิติ ซึ่งมีรายละเอียดในหนังสือ Lognormal Distributions: Theory and Applications (ค.ศ. 1988)



ภาพที่ 1-1 รูปโค้งการแจกแจงลอการมัล และการแจกแจงไวบูลล์



ภาพที่ 1-2 แสดงอัตราการชำรุดของการแจกแจงลอการมัลและไวบูลล์

จากภาพที่ 1-1¹ ถ้าพิจารณาโค้งการแจกแจงลอการิทึม (Log-normal Distribution) และการแจกแจงไวบูลล์ (Weibull Distribution) จะมีลักษณะใกล้เคียงกัน ดังนั้นในการพิจารณาจาก กราฟจึงยากที่บอกได้ว่าข้อมูลนั้นมีความใกล้เคียงการแจกแจงแบบใด และจากภาพที่ 1-2² แสดงโค้งของอัตราการชำรุด (Failure Rate) ของการแจกแจงลอการิทึมและการแจกแจงไวบูลล์ จะมีลักษณะที่แตกต่างกัน ดังนั้นถ้าเราเลือกการแจกแจงที่ไม่ถูกต้อง แล้วนำมาวิเคราะห์อัตราการชำรุด ผลที่ได้อาจจะผิดพลาดสูง

Dumonceaux และ Antle(ค.ศ. 1973) อธิบาย Ratio of Maximum Likelihoods เพื่อใช้ในการจำแนกการแจกแจงลอการิทึม (Log-normal Distribution) 2 พารามิเตอร์ เทียบกับการแจกแจงไวบูลล์ (Weibull Distribution) 2 พารามิเตอร์

ถ้ามีค่าสังเกต $x_1, x_2, \dots, x_n$ และบนพื้นฐานของการสังเกตค่าเหล่านี้ต้องการที่จะเลือกฟังก์ชันการแจกแจงลอการิทึม 2 พารามิเตอร์ (Crow, E.L. and Shimizu, K., 1988)

$$f(x; \mu, \sigma) = \frac{1}{x\sqrt{2\pi\sigma^2}} \exp\left[-(\ln x - \mu)^2 / 2\sigma^2\right], \quad x > 0, \sigma > 0$$

หรือฟังก์ชันการแจกแจงไวบูลล์ 2 พารามิเตอร์ (Coles, S.G., 1989)

$$g(x; \beta, \alpha) = \frac{\alpha}{\beta} \left(\frac{x}{\beta}\right)^{\alpha-1} \exp\left[-\left(\frac{x}{\beta}\right)^\alpha\right], \quad x > 0, \beta > 0, \alpha > 0$$

ฟังก์ชันใดฟังก์ชันหนึ่ง การทดสอบภาวะสารูปสมมติ (Goodness of Fit Test : GOF) เป็นวิธีการทางสถิติที่ใช้ทดสอบว่า ข้อมูลที่ได้มานั้นมีลักษณะการแจกแจง เป็นไปตามที่คาดหวังหรือไม่ โดยมีหลักการของสถิติทดสอบคือ การวัดระยะที่ห่างที่สุดระหว่างกราฟของ $S(x)$ และ $F_0(x)$ ซึ่ง $S(x)$ แทนฟังก์ชันการแจกแจงความถี่สัมพัทธ์สะสมของตัวอย่าง หรือความถี่สะสมที่สังเกตได้ในรูปของสัดส่วน (empirical distribution function : EDF) และ $F_0(x)$ แทนฟังก์ชันการแจกแจงความถี่สะสมภายใต้ H_0 (Hypothesized distribution function) ซึ่งมีการทดสอบหลายวิธีด้วยกัน เช่น Kolmogorov-Smirnov test, Cramer-von Mises test เป็นต้น

สำหรับการวิจัยครั้งนี้ ผู้วิจัยมีความสนใจศึกษาวิธีการทดสอบ GOF ในการจำแนกการแจกแจงลอการิทึมและการแจกแจงแบบไวบูลล์ โดยจะศึกษา Anderson Darling Test Statistic, Kuiper Test Statistic, Ratio of Maximum Likelihoods Test Statistic และ Kolmogorov -Smirnov Test Statistic โดยพิจารณาจากความสามารถในการควบคุมค่า

¹ McDonald, G.C., Vance, L.C. and Gibbons, D.I., "Some Tests for Discriminating Between Lognormal and Weibull Distributions-An Application to Emission Data." 1995; 488

² McDonald, G.C., Vance, L.C. and Gibbons, D.I., "Some Tests for Discriminating Between Lognormal and Weibull Distributions-An Application to Emission Data." 1995; 489

ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 แล้วทำการเปรียบเทียบประสิทธิภาพของวิธีการทดสอบ ด้วยค่าอำนาจของการทดสอบ

1.2 วัตถุประสงค์ของการวิจัย

การวิจัยครั้งนี้มีจุดมุ่งหมาย ดังนี้

1.2.1 เพื่อศึกษาวิธีการทดสอบภาวะสสารูปสนธิดีสำหรับการแจกแจงไวบูลล์ และการแจกแจงลอการิธึม ซึ่งสถิติทดสอบที่ใช้ในการวิจัย ได้แก่

- Anderson Darling Test Statistic
- Kuiper Test Statistic
- Kolmogorov -Smirnov Test Statistic
- Ratio of Maximum Likelihoods Test Statistic

1.2.2 เพื่อเปรียบเทียบวิธีการทดสอบภาวะสสารูปสนธิดีทั้ง 4 วิธี โดยพิจารณาจากความสามารถในการควบคุมความคลาดเคลื่อนประเภทที่ 1 (Type I error) และค่าอำนาจของการทดสอบ (Power of the test)

1.3 สมมติฐานของการวิจัย

สถานการณ์ที่ถูกกำหนดภายใต้ขอบเขตการวิจัย สถิติทดสอบ Ratio of Maximum Likelihoods Test Statistic ให้ค่าอำนาจการทดสอบสูงกว่า Anderson-Darling Test, Kolmogorov-Smirnov Test และ Kuiper Test

1.4 ขอบเขตของการวิจัย

1.4.1 กำหนดลักษณะของข้อมูล มีการแจกแจง 2 การแจกแจง ดังนี้

1.4.1.1 การแจกแจงแบบลอการิธึม (Log-normal Distribution) ที่พารามิเตอร์ ดังนี้

$$\mu = 1 \text{ และ } \sigma = 0.5$$

$$\mu = 1 \text{ และ } \sigma = 1$$

$$\mu = 1 \text{ และ } \sigma = 1.5$$

$$\mu = 1 \text{ และ } \sigma = 2$$

$$\mu = 1 \text{ และ } \sigma = 3$$

1.4.1.2 การแจกแจงแบบไวบูลล์ (Weibull Distribution) ที่พารามิเตอร์ ดังนี้

$$\beta = 1 \text{ และ } \alpha = 0.5$$

$$\beta = 1 \text{ และ } \alpha = 1$$

$$\beta = 1 \text{ และ } \alpha = 1.5$$

$$\beta = 1 \text{ และ } \alpha = 2$$

$$\beta = 1 \text{ และ } \alpha = 3$$

1.4.2 สถิติทดสอบที่ใช้ในการทดสอบสมมุติฐานมี 4 วิธี ได้แก่

1.4.2.1 สถิติทดสอบ Anderson-Darling

$$A = \left[- \sum_{i=1}^n \frac{(2i-1)}{n} \{ \ln(F_0(x_i)) + \ln(1 - F_0(x_{n+1-i})) \} \right] - n$$

1.4.2.2 สถิติทดสอบ Kolmogorov - Smirnov

$$KS = \max(D^+, D^-)$$

$$\text{เมื่อ } D^+ = \max_{i=1, \dots, n} \left\{ \frac{i}{n} - F_0(x_i) \right\}$$

$$\text{และ } D^- = \max_{i=1, \dots, n} \left\{ F_0(x_i) - \frac{i-1}{n} \right\}$$

1.4.2.3 สถิติทดสอบ Kuiper

$$V = D^+ + D^-$$

$$\text{เมื่อ } D^+ = \max_{i=1, \dots, n} \left\{ \frac{i}{n} - F_0(x_i) \right\}$$

$$\text{และ } D^- = \max_{i=1, \dots, n} \left\{ F_0(x_i) - \frac{i-1}{n} \right\}$$

1.4.2.4 สถิติทดสอบ Ratio of Maximum Likelihoods

ก) ตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงไวบูลล์

$$RML = \frac{\max_{\hat{\beta}, \hat{\alpha}} \prod_{i=1}^n f_w(x_i; \hat{\beta}, \hat{\alpha})}{\max_{\hat{\mu}, \hat{\sigma}} \prod_{i=1}^n f_{\ln}(x_i; \hat{\mu}, \hat{\sigma})}$$

$$(RML)^{\frac{1}{n}} = (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_{i=1}^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}}$$

$$\text{เมื่อ } \hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n} \text{ และ } \hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n}$$

ข) ตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงลอกนอร์มัล

$$RML = \frac{\max_{\hat{\mu}, \hat{\sigma}} \prod_{i=1}^n f_{\ln}(x_i; \hat{\mu}, \hat{\sigma})}{\max_{\hat{\beta}, \hat{\alpha}} \prod_{i=1}^n f_{\ln}(x_i; \hat{\beta}, \hat{\alpha})}$$

$$(RML)^{\frac{1}{n}} = (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_{i=1}^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}}$$

$$\text{เมื่อ } \hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n} \text{ และ } \hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n}$$

1.4.3 กำหนดระดับนัยสำคัญ (α_1) 3 ระดับ ได้แก่ 0.01, 0.05, 0.10

1.4.4 ขนาดตัวอย่างที่ใช้ในการศึกษา ได้แก่ 20, 30, 40, 50

1.4.5 การวิจัยครั้งนี้ได้จำลองข้อมูลให้มีสถานการณ์ตามที่กำหนดข้างต้น โดยใช้โปรแกรม R และทำการจำลองข้อมูลซ้ำๆ กัน 1,000 ครั้ง ในแต่ละสถานการณ์

จำนวนสถานการณ์ แบ่งออกเป็น 3 ส่วนด้วยกัน ดังนี้

1. การคำนวณหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1

$2 \times 5 \times 4 \times 3 = 120$ สถานการณ์

2. การคำนวณหาค่าอำนาจการทดสอบ $2 \times 5 \times 4 \times 3 = 120$ สถานการณ์

รวมจำนวนสถานการณ์ทั้งหมด 240 สถานการณ์

1.5 เกณฑ์ที่ใช้ในการศึกษาประสิทธิภาพของตัวแบบ

เกณฑ์การเปรียบเทียบสถิติทดสอบแต่ละวิธี สำหรับการแจกแจงไวบูลล์และการแจกแจงล็อกนอร์มัล พิจารณาเป็น 2 ขั้นตอนดังนี้

1.5.1 พิจารณาความสามารถในการควบคุมความคลาดเคลื่อนประเภทที่ 1 โดยการนำค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของการทดลอง ในแต่ละสถานการณ์มาเปรียบเทียบกับเกณฑ์ ซึ่งในการวิจัยครั้งนี้ได้แสดงเกณฑ์ที่ใช้ในการตัดสินใจ 2 เกณฑ์ด้วยกัน คือ เกณฑ์ของ Cochran และ เกณฑ์ของ Bradley รายละเอียดของเกณฑ์ทั้งสองได้เสนอไว้ในบทที่ 4

1.5.2 พิจารณาค่าอำนาจการทดสอบของสถิติทดสอบทั้ง 4 วิธี โดยพิจารณาสถานการณ์ที่ผ่านเกณฑ์จากข้อ 1 แล้ว มาหาค่าอำนาจการทดสอบ เพื่อเปรียบเทียบสถิติทดสอบทั้ง 4 วิธีว่าสถิติทดสอบวิธีใดให้ค่าอำนาจการทดสอบสูงที่สุด

1.6 ประโยชน์ของการวิจัยที่คาดว่าจะได้รับ

1.6.1 เพื่อเป็นแนวทางในการตัดสินใจในการเลือกใช้วิธีการทดสอบภาวะรูปสถิติ ในการจำแนกการแจกแจง ได้อย่างเหมาะสมและมีประสิทธิภาพ

1.6.2 เพื่อเป็นแนวทางในการตัดสินใจ ในการใช้การแจกแจงแบบล็อกนอร์มัลและการแจกแจงไวบูลล์ได้ไม่ผิดพลาด

1.6.3 เพื่อเป็นประโยชน์ต่อการประยุกต์ใช้ต่างๆ เช่น ทางด้านความเชื่อถือได้ของผลิตภัณฑ์ เป็นต้น

บทที่ 2

ทฤษฎีและตัวสถิติที่ใช้ในการวิจัย

ในบทนี้จะกล่าวถึงรายละเอียดทฤษฎีต่างๆ สถิติทดสอบที่ใช้ในการทดสอบภาวะสมมติฐานดี และการแจกแจงต่างๆ ของข้อมูลที่ใช้ในการวิจัยรวมทั้งผลงานวิจัยที่เกี่ยวข้อง

2.1 การแจกแจงของข้อมูลที่ใช้ในการวิจัย

2.1.1 การแจกแจงไวบูลล์ (Weibull Distribution) 2 พารามิเตอร์

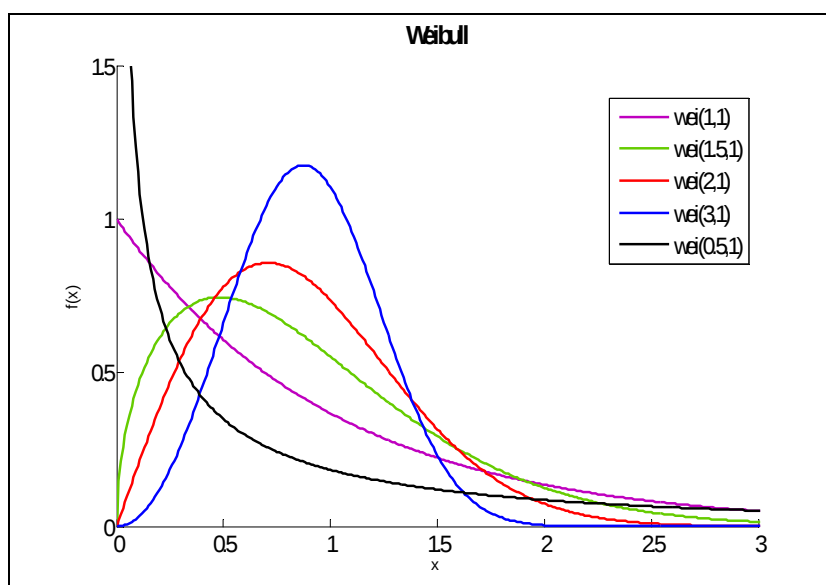
มีฟังก์ชันความหนาแน่นและฟังก์ชันสะสมของตัวแปรสุ่ม x อยู่ในรูป

$$f(x) = \begin{cases} \frac{\alpha}{\beta} \left(\frac{x}{\beta}\right)^{\alpha-1} \exp\left[-\left(\frac{x}{\beta}\right)^{\alpha}\right] & ; x > 0 \\ 0 & ; x \text{ อื่นๆ} \end{cases}$$

$$F(x) = \begin{cases} 1 - \exp\left[-\left(\frac{x}{\beta}\right)^{\alpha}\right] & ; x > 0 \\ 0 & ; x \text{ อื่นๆ} \end{cases}$$

โดยที่ α เป็นพารามิเตอร์แสดงรูปร่าง (Shape Parameter), $\alpha > 0$

และ β เป็นพารามิเตอร์แสดงขนาด (Scale Parameter), $\beta > 0$



ภาพที่ 2-1 กราฟของการแจกแจงไวบูลล์ที่ $\beta = 1$ และ $\alpha = 0.5, 1, 1.5, 2$ และ 3

2.1.2 การแจกแจงลอการิธึม (Lognormal distribution)

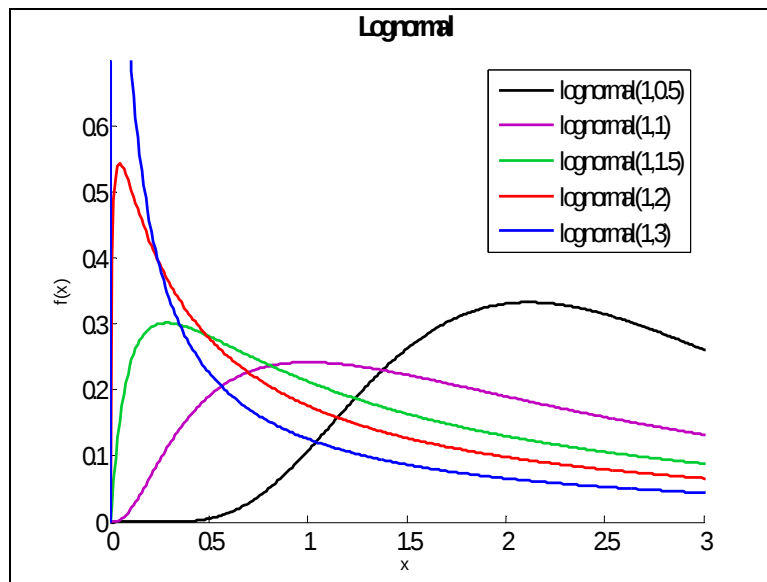
มีฟังก์ชันความหนาแน่นและฟังก์ชันสะสมของตัวแปรสุ่ม x อยู่ในรูป

$$f(x) = \begin{cases} \frac{1}{x\sqrt{2\pi\sigma^2}} \exp\left[-\frac{(\ln x - \mu)^2}{2\sigma^2}\right] & ; x > 0 \\ 0 & ; x \text{ อื่นๆ} \end{cases}$$

$$F(x) = \begin{cases} \Phi\left(\frac{\ln x - \mu}{\sigma}\right) & ; x > 0 \\ 0 & ; x \text{ อื่นๆ} \end{cases}$$

โดยที่ σ เป็นพารามิเตอร์แสดงรูปร่าง (Shape Parameter), $\sigma > 0$

และ μ เป็นพารามิเตอร์แสดงขนาด (Scale Parameter), $\mu \in (-\infty, \infty)$



ภาพที่ 2-2 กราฟของการแจกแจงลอการิธึมที่ $\mu = 1$ และ $\sigma = 0.5, 1, 1.5, 2$ และ 3

2.2 การเปรียบเทียบการแจกแจงและอัตราการชำรุดของการแจกแจงลอการิธึมและการแจกแจงไวบูลล์

ถ้าพิจารณาจากรูปที่ 2.1 และ รูปที่ 2.2 การแจกแจงลอการิธึมเป็นการแจกแจงที่มีความเบ้ และจะมีหางของการแจกแจงที่ยาวกว่าหางของการแจกแจงไวบูลล์ ขึ้นอยู่กับค่าพารามิเตอร์ จะแสดงการเปรียบเทียบค่าต่างๆ ของการแจกแจงลอการิธึมและการแจกแจงไวบูลล์ในตารางที่ 2-1

ตารางที่ 2-1 การเปรียบเทียบการแจกแบบลอการิธึมและการแจกแบบไวบูลล์

Distribution	Lognormal(μ, σ)	Weibull(α, β)
Mean	$\exp\left(\mu + \frac{\sigma^2}{2}\right)$	$\beta \Gamma\left(1 + \frac{1}{\alpha}\right)$
Median	$\exp(\mu)$	$\beta \left[(\ln 2)^{\frac{1}{\alpha}} \right]$
Mode	$\exp(\mu - \sigma^2)$	$\beta \left(1 - \frac{1}{\alpha}\right)^{1/\alpha}$
Variance	$\exp(2\mu + \sigma^2) [\exp(\sigma^2) - 1]$	$\beta^2 \left[\Gamma\left(1 + \frac{2}{\alpha}\right) - \Gamma^2\left(1 + \frac{1}{\alpha}\right) \right]$
Percentile	$\exp(\mu + v_q \sigma)$	$\beta \left[(\ln(1-q)^{-1})^{\frac{1}{\alpha}} \right]$

ให้ v_q เป็นค่าของ Quantile ที่ q ของการแจกแจงปกติ และ $\Gamma(x) = \int_0^\infty z^{x-1} \exp(-z) dz$

การเปรียบเทียบส่วนหางของการแจกแบบลอการิธึมและการแจกแบบไวบูลล์สามารถทำได้โดยเทียบค่า Median และค่า quantile ที่ .95 ($q_{0.95}$)

กำหนดให้ $\mu = 1$ และ $\sigma = 1$ จากตารางที่ 2-1 จะได้ว่า Median และค่า quantile ที่ .95 ของการแจกแจงลอการิธึม คือ

$$\text{Median} = \exp(1)$$

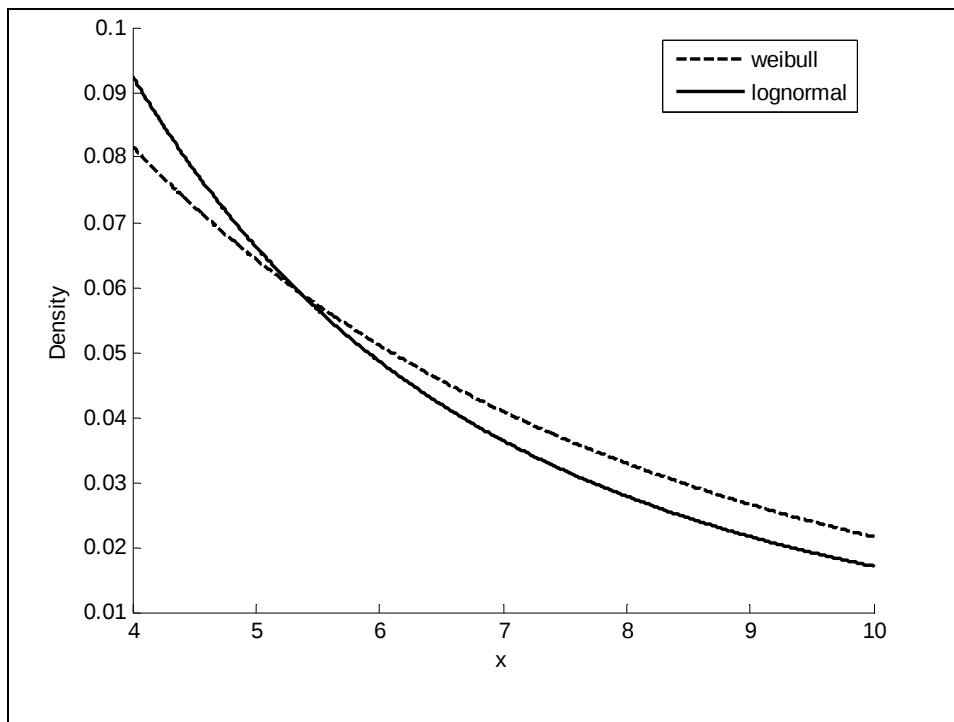
$$\begin{aligned} q_{0.95} &= \exp(1 + v_{0.95}(1)) \\ &= \exp(2.645) \end{aligned}$$

ดังนั้นจะหาค่าพารามิเตอร์ β และ α จากการแก้สมการต่อไปนี้

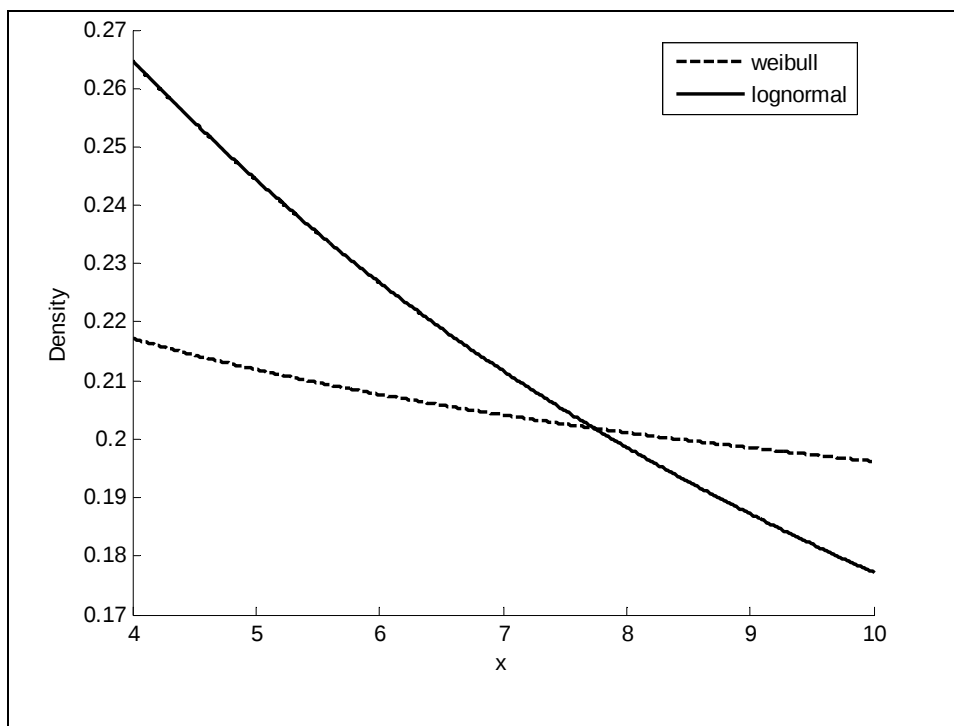
$$\begin{aligned} \exp(1) &= \beta \left[(\ln 2)^{\frac{1}{\alpha}} \right] \\ \exp(2.645) &= \beta \left[(\ln(1-0.95)^{-1})^{\frac{1}{\alpha}} \right] \end{aligned}$$

ซึ่งจากการแก้สมการได้ค่า พารามิเตอร์ $\beta = 4.104$ และ $\alpha = 0.889$

จากภาพที่ 2-3 แสดงการเปรียบเทียบการแจกแจงลอการิธึมที่พารามิเตอร์ $\mu = 1$ และ $\sigma = 1$ กับ การแจกแจงไวบูลล์ที่พารามิเตอร์ $\beta = 4.104$ และ $\alpha = 0.889$ จะเห็นได้ว่าลักษณะของโค้งการแจกแจงทั้งสองมีความใกล้เคียงกันมาก ดังนั้นในการพิจารณาจาก กราฟจึงยากที่บอกได้ว่าข้อมูลนั้นมีความใกล้เคียงการแจกแจงแบบใด



ภาพที่ 2-3 รูปโค้งการแจกแจงลอการิทึมที่ $\mu = 1$ และ $\sigma = 1$
และการแจกแจงไวบูลล์ที่ $\beta = 4.104$ และ $\alpha = 0.889$



ภาพที่ 2-4 แสดงอัตราการชำรุดของการแจกแจงลอการิทึมที่ $\mu = 1$ และ $\sigma = 1$
และการแจกแจงไวบูลล์ที่ $\beta = 4.104$ และ $\alpha = 0.889$

ถ้าเราพิจารณาอัตราการชำรุด (Failure Rate) หรือแทนด้วย $r(t)$ (Rigdon, S.E. and Basu,2000) ซึ่งเราสามารถคำนวณค่าอัตราการชำรุดได้จาก

$$r(t) = f(t) / [1 - F(t)]$$

เมื่อ $f(t)$ เป็นฟังก์ชันการแจกแจง และ $F(t)$ เป็นฟังก์ชันการแจกแจงสะสม

จากภาพที่ 2-4 แสดงการเปรียบเทียบอัตราการชำรุดของการแจกแจงลอการิธึมที่พารามิเตอร์ $\mu=1$ และ $\sigma=1$ กับ การแจกแจงไวบูลล์ที่พารามิเตอร์ $\beta=4.104$ และ $\alpha=0.889$ จะมีลักษณะที่แตกต่างกัน ดังนั้นถ้าเราเลือกการแจกแจงที่ไม่ถูกต้อง แล้วนำมาวิเคราะห์อัตราการชำรุด ผลที่ได้อาจจะผิดพลาดสูง

2.3 การทดสอบภาวะสารูปสนิทธิ

การทดสอบภาวะสารูปสนิทธิ (Goodness-of-Fit Test:GOF) คือกระบวนการทางสถิติที่ใช้ทดสอบสมมติฐานเพื่อยอมรับ หรือปฏิเสธการแจกแจงความน่าจะเป็นที่แสดงถึงลักษณะของประชากร ซึ่งสมมติฐานที่ใช้ในการทดสอบคือ

$$H_0 : S(x) = F_0(x)$$

หรือไม่มีความแตกต่างกันระหว่างการแจกแจงของข้อมูลที่สังเกตได้กับการแจกแจงที่คาดหวังไว้

$$H_1 : S(x) \neq F_0(x)$$

หรือมีความแตกต่างกันระหว่างการแจกแจงของข้อมูลที่สังเกตได้กับการแจกแจงที่คาดหวังไว้

ให้ $x_1, x_2, \dots, x_n$ เป็นตัวอย่างสุ่มขนาด n ที่มาจากประชากรที่มีฟังก์ชันการแจกแจง $S(x)$ และ $S(x)$ เป็นฟังก์ชันการแจกแจงของข้อมูลที่สังเกตได้ และ $F_0(x)$ เป็นฟังก์ชันการแจกแจงประชากรที่กำหนด

2.4 ตัวสถิติทดสอบที่ใช้ในการวิจัย

สถิติทดสอบที่ใช้ในการทดสอบ Goodness-of-fit test (GOF) ในการวิจัยครั้งนี้มี 4 วิธี ซึ่งมีรายละเอียดดังนี้

2.4.1 Anderson Darling Test Statistic

Anderson Darling (AD) Test (Stephens, 1974) เป็นสถิติทดสอบเพื่อใช้ทดสอบลักษณะของประชากรว่าเป็นการแจกแจงที่คาดหวังไว้หรือไม่ซึ่ง สถิติทดสอบAD เป็นวิธีที่ได้มาจากการดัดแปลงจาก Kolmogorov -Smirnov (KS)Test

สถิติทดสอบ KSใช้สำหรับการแจกแจงได้ทุกการแจกแจงโดยที่ค่าวิกฤต (critical values) สามารถหาได้โดยที่ไม่จำเป็นต้องระบุการแจกแจงที่จะใช้ทดสอบ แต่สถิติทดสอบ AD จำเป็นต้องระบุการแจกแจงที่จะใช้ทดสอบ ในการคำนวณค่าวิกฤต AD มีข้อดีคือเป็นการทดสอบที่มีความไวในการคำนวณ แต่ AD ก็มีข้อเสียคือ ในการคำนวณค่าวิกฤตจะต้อง

แยกแต่ละการแจกแจง แต่ในปัจจุบันตารางค่าวิกฤต ของ การแจกแจงนอร์มัล (Normal distribution) การแจกแจงล็อกนอร์มัล (Lognormal distribution) การแจกแจงล็อกนอร์มัล (Lognormal distribution) การแจกแจงโลจิสติก (Logistic distribution) หาได้ทั่วไป

การทดสอบ AD คำนวณสถิติทดสอบดังนี้

$$A = \left[- \sum_{i=1}^n \frac{(2i-1)}{n} \{ \ln(F_0(x_i)) + \ln(1 - F_0(x_{n+1-i})) \} \right] - n$$

เมื่อ $F_0[x_i]$ เป็นฟังก์ชันการแจกแจงสะสมตามการแจกแจงที่คาดหวังไว้ ของ x_i ซึ่ง x_i เป็นของมูลที่เรียงลำดับไว้

ในการวิจัยครั้งนี้ผู้วิจัยจะเลือก AD เพื่อเป็นตัวสถิติทดสอบในการเปรียบเทียบการแจกแจงล็อกนอร์มัล และการแจกแจงไวบูลล์ ดังนี้

2.4.1.1 การใช้ AD ในกรณีที่ให้ H_0 : การแจกแจงล็อกนอร์มัล

และ H_1 : การแจกแจงไวบูลล์

เริ่มจากการประมาณค่าพารามิเตอร์ของการแจกแจงล็อกนอร์มัล คือ σ (scale parameter) และ μ (shape parameter) คำนวณค่า A โดยให้ $F_0[x_i]$ คือ ฟังก์ชันการแจกแจงสะสมของการแจกแจงล็อกนอร์มัลตามพารามิเตอร์ที่ได้จากการประมาณค่าจะปฏิเสธสมมติฐาน H_0 : การแจกแจงล็อกนอร์มัล และยอมรับ H_1 : การแจกแจงไวบูลล์ เมื่อ $A^* \geq A_c$ ซึ่ง A_c คือ ค่าวิกฤตของ AD โดยจะแสดงค่า A_c ในตารางที่ 2-2 โดย D'Agostino and Stephens (1986) คำนวณค่า A^* ดังนี้

$$A^* = A \left(1 + \frac{0.75}{n} + \frac{2.25}{n^2} \right)$$

ตารางที่ 2-2 ค่า A_c สำหรับ H_0 : การแจกแจงล็อกนอร์มัล และ H_1 : การแจกแจงไวบูลล์

α_1	0.10	0.05	0.01
A_c	0.631	0.752	1.035

ตารางอ้างอิงมาจาก D'Agostino and Stephens, 1986, Table 4.7, p.123

2.4.1.2 การใช้ AD ในกรณีที่ให้ H_0 : การแจกแจงไวบูลล์

และ H_1 : การแจกแจงล็อกนอร์มัล

เริ่มจากการประมาณค่าพารามิเตอร์ของการแจกแจงไวบูลล์ คือ β (scale parameter) และ α (shape parameter) คำนวณค่า A โดยให้ $F_0[x_i]$ คือ ฟังก์ชันการแจกแจงสะสมของการแจกแจงไวบูลล์ตามพารามิเตอร์ที่ได้จากการประมาณค่า จะปฏิเสธสมมติฐาน H_0 : การแจกแจงไวบูลล์ และยอมรับ H_1 : การแจกแจงล็อกนอร์มัล เมื่อ $A^* \geq A_c$ ซึ่ง A_c

คือ critical values ของ AD โดยจะแสดงค่า $A^* \geq A_c$ ซึ่ง A_c ในตารางที่ 2-3 โดย D'Agostino and Stephens (1986) คำนวณค่า A^* ดังนี้

$$A^* = A \left(1 + \frac{0.2}{\sqrt{n}} \right)$$

ตารางที่ 2-3 ค่า A_c สำหรับ H_0 : การแจกแจงไวบูลล์ และ H_1 : การแจกแจงลอกนอร์มัล

α_1	0.10	0.05	0.01
A_c	0.637	0.757	1.038

ตารางอ้างอิงมาจาก D'Agostino and Stephens, 1986, Table 4.17, p.146

จะแสดงตัวอย่างการใช้สถิติทดสอบ AD จากตัวอย่างที่ 1 ดังต่อไปนี้

ตัวอย่างที่ 1 สร้างข้อมูลจากการแจกแจงไวบูลล์ ($\alpha = 10, \beta = 2$) ดังนี้

11.722 10.429 8.020 7.578 1.430 4.115

จากข้อมูลทั้ง 6 คำนวณค่าประมาณ $\beta = 2.12$ และประมาณค่า $\alpha = 8.1$

i	x_i	$F_0[x_i]$	$\ln(F_0[x_i])$	$\ln(1 - F_0[x_{n+1-i}])$	$((2i-1)/n) \{ \ln(F_0[x_i]) + \ln(1 - F_0[x_{n+1-i}]) \}$
1	1.43	0.0251	-3.687	-2.185	-0.979
2	4.115	0.212	-1.552	-1.706	-1.629
3	7.578	0.580	-0.544	-0.978	-1.269
4	8.02	0.624	-0.472	-0.867	-1.562
5	10.429	0.818	-0.200	-0.238	-0.658
6	11.722	0.887	-0.119	-0.0254	-0.265

$$\sum_{i=1}^n \frac{(2i-1)}{n} \{ \ln(F_0[x_i]) + \ln(1 - F_0[x_{n+1-i}]) \} = -6.362$$

$$\text{จะได้ } A = -(-6.362) - 6$$

$$= 0.362$$

$$\text{และ } A^* = \left(1 + \frac{0.2}{\sqrt{6}} \right) 0.362$$

$$= 0.391$$

จากตารางที่ 2-3 จะเห็นว่าทุกระดับนัยสำคัญ (α_1) ค่า $A^* < A_c$

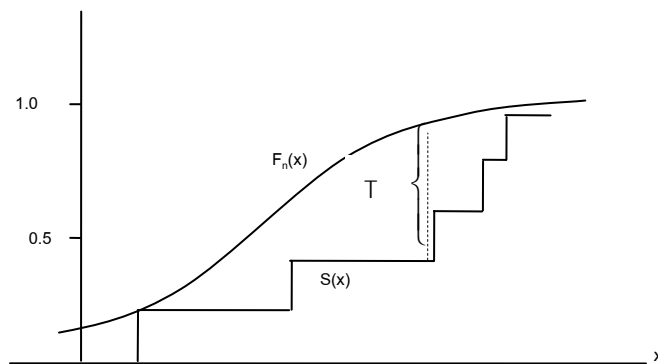
ดังนั้น จะยอมรับสมมติฐาน H_0 กล่าวคือ ข้อมูลชุดนี้มีการแจกแจงแบบไวบูลล์

2.4.2 Kolmogorov -Smirnov Test Statistic

Lilliefors, H.W. (1969) ได้เสนอตัวสถิติทดสอบ Kolmogorov -Smirnov (KS) เพื่อใช้ทดสอบการแจกแจงแบบเอ็กโปเนนเชียล ซึ่งสถิติทดสอบ KS Test เป็นสถิติทดสอบที่ใช้หลักการของการวัดระยะที่ห่างที่สุดระหว่างกราฟของ $S(x_i)$ และ $F_0[x_i]$ โดยที่ $S(x_i)$ แทนฟังก์ชันการแจกแจงความถี่สัมพัทธ์สะสมของตัวอย่าง หรือความถี่สะสมที่สังเกตได้ในรูปของสัดส่วน (empirical distribution function)

$$S(x_i) = \frac{n(i)}{N}$$

เมื่อ $n(i)$ คือ จำนวนข้อมูลที่มีค่าน้อยกว่าหรือเท่ากับ x_i , N คือ จำนวนข้อมูลทั้งหมด และ $F_0[x_i]$ แทนฟังก์ชันการแจกแจงความถี่สะสมภายใต้ H_0 (hypothesized distribution function) ดังแสดงในภาพที่ 2-5



ภาพที่ 2-5 แสดงฟังก์ชันการแจกแจงความถี่สัมพัทธ์สะสมของตัวอย่าง หรือความถี่สะสมที่สังเกตได้ในรูปของสัดส่วน: $S(x_i)$, ฟังก์ชันการแจกแจงความถี่สะสมภายใต้ H_0 : $F_0[x_i]$ และค่าสถิติ KS : T

โดย Bain, L.J. and Engelhardt(1983) คำนวณค่าสถิติทดสอบดังนี้

$$KS = \max(D^+, D^-)$$

$$\text{เมื่อ } D^+ = \max_{i=1, \dots, n} \left\{ \frac{i}{n} - F_0(x_i) \right\}$$

$$\text{และ } D^- = \max_{i=1, \dots, n} \left\{ F_0(x_i) - \frac{i-1}{n} \right\}$$

เมื่อ $F_0[x_i]$ เป็นฟังก์ชันการแจกแจงสะสมตามการแจกแจงที่คาดหวัง ของ x_i ซึ่ง x_i เป็นของข้อมูลที่เรียงลำดับ

ในการวิจัยครั้งนี้ผู้วิจัยจะเลือก KS เพื่อเป็นตัวสถิติทดสอบในการเปรียบเทียบการแจกแจงลอการิธึม และการแจกแจงไวบูลล์ ดังนี้

2.4.2.1 การใช้ KS ในกรณีที่ให้ H_0 : การแจกแจงลอกลอนอร์มัล

และ H_1 : การแจกแจงไวบูลล์

เริ่มจากการประมาณค่าพารามิเตอร์ของการแจกแจงลอกลอนอร์มัล คือ σ (scale parameter) และ μ (shape parameter) คำนวณค่า D โดยให้ $F_0[x_i]$ คือ ฟังก์ชันการแจกแจงสะสมของการแจกแจงลอกลอนอร์มัลตามพารามิเตอร์ที่ได้จากการประมาณค่าจะปฏิเสธสมมติฐาน H_0 : การแจกแจงลอกลอนอร์มัล และยอมรับ H_1 : การแจกแจงไวบูลล์ เมื่อ $D^* \geq D_c$ ซึ่ง D_c คือ ค่าวิกฤตของ KS โดยจะแสดงค่า D_c ในตารางที่ 2-4 โดย Bain, L.J. and Engelhardt (1973) คำนวณค่า D^* ดังนี้

$$D^* = D \left(\sqrt{n} - 0.1 + \frac{0.85}{\sqrt{n}} \right)$$

ตารางที่ 2-4 ค่า D_c สำหรับ H_0 : การแจกแจงลอกลอนอร์มัล และ H_1 : การแจกแจงไวบูลล์

α_1	0.10	0.05	0.01
D_c	0.819	0.895	1.035

ตารางอ้างอิงมาจาก Bain, L.J. and Engelhardt, M., 1973, Table 11, p.613

2.4.2.2 การใช้ KS ในกรณีที่ให้ H_0 : การแจกแจงไวบูลล์

และ H_1 : การแจกแจงลอกลอนอร์มัล

เริ่มจากการประมาณค่าพารามิเตอร์ของการแจกแจงไวบูลล์ คือ β (scale parameter) และ α (shape parameter) คำนวณค่า D โดยให้ $F_0[x_i]$ คือ ฟังก์ชันการแจกแจงสะสมของการแจกแจงไวบูลล์ตามพารามิเตอร์ที่ได้จากการประมาณค่า จะปฏิเสธสมมติฐาน H_0 : การแจกแจงไวบูลล์ และยอมรับ H_1 : การแจกแจงลอกลอนอร์มัล เมื่อ $D^* \geq D_c$ ซึ่ง D_c คือ critical values ของ KS โดยจะแสดงค่า D_c ในตารางที่ 2-5 โดย Bain, L.J. and Engelhardt (1973) คำนวณค่า D^* ดังนี้

$$D^* = (\sqrt{n})D$$

ตารางที่ 2-5 ค่า D_c สำหรับ H_0 : การแจกแจงไวบูลล์ และ H_1 : การแจกแจงลอกลอนอร์มัล

α_1	0.10	0.05	0.01
D_c	0.803	0.874	1.007

ตารางอ้างอิงมาจาก Bain, L.J. and Engelhardt, M., 1973, Table 11, p.613

จะแสดงตัวอย่างการใช้สถิติทดสอบ KS จากตัวอย่างที่ 1 ดังต่อไปนี้

11.722 10.429 8.020 7.578 1.430 4.115

จากข้อมูลทั้ง 6 คำนวณค่าประมาณ $\beta=2.12$ และประมาณค่า $\alpha=8.1$

i	x_i	$F_0[x_i]$	$\frac{i}{N} - F_0(x_i)$	$F_0(x_i) - \frac{i-1}{N}$
1	1.43	0.0251	0.142	0.0251
2	4.115	0.212	0.121	0.045
3	7.578	0.580	-0.08	0.247
4	8.02	0.624	0.043	0.124
5	10.429	0.818	0.015	0.151
6	11.722	0.887	0.113	0.054

$$\text{จาก } D_N^+ = 0.142$$

$$\text{และ } D_N^- = 0.247$$

$$\text{ดังนั้น } D = \max(D_N^+, D_N^-)$$

$$= 0.247$$

$$D^* = (\sqrt{6})0.247$$

$$= 0.605$$

จากตารางที่ 2-5 จะเห็นว่าทุกระดับนัยสำคัญ (α_1) ค่า $D^* < D_c$

ดังนั้น จะยอมรับสมมติฐาน H_0 กล่าวคือ ข้อมูลชุดนี้มีการแจกแจงแบบไวบูลล์

2.4.3 Kuiper Test Statistic

Berr, D. R. and Shudde, R. H. (1973) ได้เสนอตัวสถิติทดสอบ kuiper(K) เพื่อใช้ทดสอบการแจกแจงแบบยูนิฟอร์ม ซึ่งสถิติทดสอบ K เป็นสถิติที่ปรับปรุงมาจากสถิติทดสอบ KS เพื่อใช้ในการทดสอบลักษณะของประชากรว่าเป็นการแจกแจงที่คาดหวังไว้หรือไม่ สถิติ K นี้มีหลักการทดสอบเช่นเดียวกับ KS คือ การวัดระยะที่ห่างที่สุดระหว่างกราฟของ $S(x_i)$ และ $F_0[x_i]$ โดยที่ $S(x_i)$ แทนฟังก์ชันการแจกแจงความถี่สัมพัทธ์สะสมของตัวอย่าง หรือความถี่สะสมที่สังเกตได้ในรูปของสัดส่วน (empirical distribution function) และ $F_0[x_i]$ แทนฟังก์ชันการแจกแจงความถี่สะสมภายใต้ H_0 (hypothesized distribution function) ของ D_N^+ และ D_N^- แล้วนำค่าที่มากที่สุดของทั้ง 2 มารวมกัน คำนวณค่าสถิติทดสอบ ดังนี้

คำนวณค่าสถิติทดสอบดังนี้

$$V = D^+ + D^-$$

$$\text{เมื่อ } D^+ = \max_{i=1, \dots, n} \left\{ \frac{i}{n} - F_0(x_i) \right\}$$

$$\text{และ } D^- = \max_{i=1, \dots, n} \left\{ F_0(x_i) - \frac{i-1}{n} \right\}$$

เมื่อ $F_0[x_i]$ เป็นฟังก์ชันการแจกแจงสะสมตามการแจกแจงที่คาดหวังไว้ ของ x_i ซึ่ง x_i เป็นของมูลที่เรียงลำดับไว้

ในการวิจัยครั้งนี้ผู้วิจัยจะเลือก K เพื่อเป็นตัวสถิติทดสอบในการเปรียบเทียบการแจกแจงลอการิธึม และการแจกแจงไวบูลล์ ดังนี้

2.4.3.1 การใช้ K ในกรณีที่ให้ H_0 : การแจกแจงลอการิธึม

และ H_1 : การแจกแจงไวบูลล์

เริ่มจากการประมาณค่าพารามิเตอร์ของการแจกแจงลอการิธึม คือ σ (scale parameter) และ μ (shape parameter) คำนวณค่า V โดยให้ $F_0[x_i]$ คือ ฟังก์ชันการแจกแจงสะสมของการแจกแจงลอการิธึมตามพารามิเตอร์ที่ได้จากการประมาณค่า จะปฏิเสธสมมติฐาน H_0 : การแจกแจงลอการิธึม และยอมรับ H_1 : การแจกแจงไวบูลล์ เมื่อ $V^* \geq V_c$ ซึ่ง V_c คือค่าวิกฤตของ K โดยจะแสดงค่า V_c ในตารางที่ 2-6 โดย Bain, L.J. and Engelhardt(1973) คำนวณค่า V^* ดังนี้

$$V^* = V \left(\sqrt{n} + 0.05 + \frac{0.82}{\sqrt{n}} \right)$$

ตารางที่ 2-6 ค่า V_c สำหรับ H_0 : การแจกแจงลอการิธึม และ H_1 : การแจกแจงไวบูลล์

α_1	0.10	0.05	0.01
V_c	1.386	1.489	1.693

ตารางอ้างอิงมาจาก Bain, L.J. and Engelhardt, M., 1973, Table 11, p.613

2.4.3.2 การใช้ K ในกรณีที่ให้ H_0 : การแจกแจงไวบูลล์

และ H_1 : การแจกแจงลอการิธึม

เริ่มจากการประมาณค่าพารามิเตอร์ของการแจกแจงไวบูลล์ คือ β (scale parameter) และ α (shape parameter) คำนวณค่า V โดยให้ $F_0[x_i]$ คือ ฟังก์ชันการแจกแจงสะสมของการแจกแจงไวบูลล์ตามพารามิเตอร์ที่ได้จากการประมาณค่า จะปฏิเสธสมมติฐาน H_0 : การแจกแจงไวบูลล์ และยอมรับ H_1 : การแจกแจงลอการิธึมเมื่อ $V^* \geq V_c$ ซึ่ง V_c คือ ค่าวิกฤตของ K โดยจะแสดงค่า V_c ในตารางที่ 2-7 โดย Bain, L.J. and Engelhardt (1973) คำนวณค่า V^* ดังนี้

$$V^* = V(\sqrt{n})$$

ตารางที่ 2-7 ค่า V_c สำหรับ H_0 : การแจกแจงไวบูลล์ และ H_1 : การแจกแจงลอกนอร์มัล

α_1	0.10	0.05	0.01
V_c	1.372	1.477	1.671

ตารางอ้างอิงมาจาก Bain, L.J. and Engelhardt, M., 1973, Table 11, p.613

จะแสดงตัวอย่างการใช้สถิติทดสอบ K จากตัวอย่างที่ 1 ดังต่อไปนี้

11.722 10.429 8.020 7.578 1.430 4.115

จากข้อมูลทั้ง 6 คำนวณค่าประมาณ $\beta = 2.12$ และประมาณค่า $\alpha = 8.1$

$$\text{จาก } D_N^+ = 0.142$$

$$\text{และ } D_N^- = 0.247$$

$$\text{ดังนั้น } V = D_N^+ + D_N^-$$

$$= 0.389$$

$$V^* = (\sqrt{6})0.389$$

$$= 0.953$$

จากตารางที่ 2-7 จะเห็นว่าทุกระดับนัยสำคัญ (α_1) ค่า $V^* < V_c$

ดังนั้น จะยอมรับสมมติฐาน H_0 กล่าวคือ ข้อมูลชุดนี้มีการแจกแจงแบบไวบูลล์

2.4.4 Ratio of Maximum Likelihoods Test Statistic

ปัญหาในการเลือกการแจกแจง 2 การแจกแจงใดๆ ที่ไม่ทราบค่าพารามิเตอร์ Dumonceaux, Antle and Hass (1973) ได้เสนอวิธีการเลือกการแจกแจงของสองตัวแบบที่มีพารามิเตอร์ 2 ตัว คือ Ratio of Maximum Likelihoods (RML) Test Statistic ซึ่งเป็นวิธีที่ไม่ได้ขึ้นอยู่กับค่าพารามิเตอร์ RML มีวิธีการหาค่า critical values ที่เหมาะสมสำหรับการทดสอบนี้ โดยการสร้างข้อมูลที่ไม่สามารถคำนวณการวิเคราะห์ เราสามารถใช้ RML ในการตรวจสอบการแจกแจงลอกนอร์มัลและ การแจกแจงไวบูลล์ ดังนั้นในการวิจัยครั้งนี้ผู้วิจัยจะเลือก RML เพื่อเป็นตัวสถิติทดสอบในการเปรียบเทียบการแจกแจงลอกนอร์มัล และการแจกแจงไวบูลล์ ดังนี้

2.4.4.1 การใช้ RML ในกรณีที่ให้ H_0 : การแจกแจงลอกนอร์มัล

และ H_1 : การแจกแจงไวบูลล์

ให้ $f_{\ln}(x_i; \hat{\mu}, \hat{\sigma})$ แทนฟังก์ชันการแจกแจงลอกนอร์มัล ที่มี $\hat{\sigma}$ เป็น scale parameter และ $\hat{\mu}$ เป็น shape parameter และให้ $f_w(x_i; \hat{\beta}, \hat{\alpha})$ แทนฟังก์ชันการแจกแจงไวบูลล์ ที่มี $\hat{\beta}$ เป็น scale parameter และ $\hat{\alpha}$ เป็น shape parameter ซึ่งสามารถหาค่า $\hat{\sigma}$ และ $\hat{\mu}$ จากวิธีการประมาณค่าพารามิเตอร์ของการแจกแจงลอกนอร์มัล ซึ่งสามารถหาค่า $\hat{\beta}$ และ $\hat{\alpha}$ จากวิธีการ

ประมาณค่าพารามิเตอร์ของการแจกแจงไวบูลล์ โดย Dumonceaux, R. and Antle, C.E. (1973)
คำนวณค่าสถิติทดสอบ ดังนี้

$$RML = \frac{\max_{\hat{\beta}, \hat{\alpha}} \prod_{i=1}^n f_w(x_i; \hat{\beta}, \hat{\alpha})}{\max_{\hat{\mu}, \hat{\sigma}} \prod_{i=1}^n f_{ln}(x_i; \hat{\mu}, \hat{\sigma})}$$

$$\text{ดังนั้น} \quad (RML)^{\frac{1}{n}} = (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_{i=1}^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}}$$

$$\text{เมื่อ} \quad \hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n} \quad \text{และ} \quad \hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n}$$

จะปฏิเสธสมมติฐาน H_0 : การแจกแจงลอกนอร์มัล และยอมรับ H_1 : การแจกแจงไวบูลล์
เมื่อ $(RML)^{\frac{1}{n}} \geq (RML)_c^{\frac{1}{n}}$ ซึ่ง $(RML)_c^{\frac{1}{n}}$ คือ ค่าวิกฤตของ $(RML)^{\frac{1}{n}}$ โดยจะแสดง
ค่า $(RML)_c^{\frac{1}{n}}$ ในตารางที่ 2-8

ตารางที่ 2-8 ค่า $(RML)_c^{\frac{1}{n}}$ สำหรับ H_0 : การแจกแจงลอกนอร์มัล
และ H_1 : การแจกแจงไวบูลล์

n	$\alpha_1 = 0.10$	$\alpha_1 = 0.05$	$\alpha_1 = 0.01$
20	1.038	1.082	1.144
30	1.020	1.044	1.095
40	1.007	1.028	1.070
50	0.998	1.014	1.054

ตารางอ้างอิงมาจาก Dumonceaux, R. and Antle, C.E., 1973, Table 1, p.924

2.4.4.2 การใช้ RML ในกรณีที่ให้ H_0 : การแจกแจงไวบูลล์

และ H_1 : การแจกแจงลอกนอร์มัล

โดย Dumonceaux, R. and Antle, C.E. (1973) คำนวณค่าสถิติทดสอบ ดังนี้

$$RML = \frac{\max_{\hat{\mu}, \hat{\sigma}} \prod_{i=1}^n f_{ln}(x_i; \hat{\mu}, \hat{\sigma})}{\max_{\hat{\beta}, \hat{\alpha}} \prod_{i=1}^n f_{ln}(x_i; \hat{\beta}, \hat{\alpha})}$$

$$\text{ดังนั้น} \quad (RML)^{\frac{1}{n}} = (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_{i=1}^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}}$$

$$\text{เมื่อ } \hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n} \quad \text{และ} \quad \hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n}$$

จะปฏิเสธสมมติฐาน H_0 : การแจกแจงไวบูลล์ และยอมรับ H_1 : การแจกแจงลอกนอร์มัล

เมื่อ $(RML)^{\frac{1}{n}} \geq (RML)_c^{\frac{1}{n}}$ ซึ่ง $(RML)_c^{\frac{1}{n}}$ คือ ค่าวิกฤตของ $(RML)^{\frac{1}{n}}$ โดยจะแสดงค่า $(RML)_c^{\frac{1}{n}}$ ในตารางที่ 2-9

ตารางที่ 2-9 ค่า $(RML)_c^{\frac{1}{n}}$ สำหรับ H_0 : การแจกแจงไวบูลล์
และ H_1 : การแจกแจงลอกนอร์มัล

n	$\alpha_1 = 0.10$	$\alpha_1 = 0.05$	$\alpha_1 = 0.01$
20	1.041	1.067	1.120
30	1.019	1.041	1.088
40	1.005	1.026	1.063
50	0.995	1.016	1.045

ตารางอ้างอิงมาจาก Dumonceaux, R. and Antle, C.E., 1973, Table 2, p.925

จะแสดงตัวอย่างการใช้สถิติทดสอบ RML จากตัวอย่างที่ 2 ดังต่อไปนี้

ตัวอย่างที่ 2 อายุการใช้งานของตลับลูกปืน 20 ลูก จาก Dumonceaux, R. and Antle, C.E.(1973)

17.88	28.92	33.00	41.25	42.12	45.60
48.48	51.84	51.96	54.12	55.56	67.80
68.64	68.64	68.88	84.12	93.12	98.64
105.12	105.84				

กำหนดให้ H_0 : การแจกแจงลอกนอร์มัล

H_1 : การแจกแจงไวบูลล์

จากข้อมูลใช้โปรแกรม R คำนวณค่าประมาณ $\hat{\beta} = 69.384$ และ $\hat{\alpha} = 2.708$

$$\text{คำนวณค่า } \hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n} = \frac{80.59458}{20} \approx 4.03$$

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n} = \frac{3.978502}{20} \approx 0.20$$

$$\begin{aligned}\text{จะได้ } (RML)^{\frac{1}{n}} &= (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_1^n x_i f_w(x_i; \hat{\beta}, \hat{\alpha}) \right]^{\frac{1}{n}} \\ &= 1.039979\end{aligned}$$

จากตารางที่ 2-8 จะเห็นว่าระดับนัยสำคัญ (α_1) ที่ 0.1 ค่า $(RML)^{\frac{1}{n}} \geq (RML)_c^{\frac{1}{n}}$ ดังนั้น จะปฏิเสธสมมติฐาน H_0 กล่าวคือ ข้อมูลชุดนี้มีการแจกแจงแบบไวบูลล์ระดับนัยสำคัญ (α_1) ที่ 0.1

จากตัวอย่างเดียวกันนี้ เมื่อทดสอบใช้ RML ในกรณี

กำหนดให้ H_0 : การแจกแจงไวบูลล์

H_1 : การแจกแจงลอการิธึม

$$\begin{aligned}\text{จะได้ } (RML)^{\frac{1}{n}} &= (2\pi e \hat{\sigma}^2)^{-\frac{1}{2}} \left[\prod_1^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}} \\ &= 0.9615578\end{aligned}$$

จากตารางที่ 2-9 จะเห็นว่าทุกระดับนัยสำคัญ (α_1) ค่า $(RML)^{\frac{1}{n}} < (RML)_c^{\frac{1}{n}}$ ดังนั้น จะยอมรับสมมติฐาน H_0 กล่าวคือ ข้อมูลชุดนี้มีการแจกแจงแบบไวบูลล์ ทุกระดับนัยสำคัญ

2.5 ผลงานวิจัยที่เกี่ยวข้อง

ผลงานวิจัยของนักสถิติที่สนใจศึกษาเปรียบเทียบสถิติทดสอบ GOF มีมากมาย ในที่นี้จะเสนอผลงานวิจัยบางผลงานวิจัยเท่านั้น

ศิริรัตน์ (2539) ได้ศึกษาเปรียบเทียบค่าอำนาจการทดสอบของสถิติทดสอบในการทดสอบการแจกแจงไวบูลล์ 2 พารามิเตอร์ ($\alpha = 3, \beta = 1$) และการแจกแจง คอมเพิร์ตซ์ ($B = 0.02, c = 20$) ด้วยการทดสอบเทียบความกลมกลืนเมื่อข้อมูลถูกตัดทิ้งประเภทที่สอง จำนวนมาก สถิติทดสอบที่ใช้ได้แก่ สถิติทดสอบ KS (Kolmogorov-Sminov Test Statistic), K (Kuiper Test Statistic) และ CVM (Cramer-von Mises Test Statistic) ภายใต้ระดับนัยสำคัญสำคัญ (α) 0.25, 0.20, 0.15, 0.10, 0.05 และ 0.01 ที่ขนาดตัวอย่าง 100, 300, 500 และ 700 เปอร์เซนต์การถูกตัดทิ้ง 90%, 95% และ 99% ยกเว้นที่ขนาดตัวอย่าง 100 และ 300 มีเปอร์เซนต์การถูกตัดทิ้ง 90% และ 95% ได้ผลสรุปดังนี้

การแจกแจงไวบูลล์ ในทุกกรณีศึกษาสามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้ด้วยเกณฑ์ของ Cochran และเมื่อพิจารณาอำนาจการทดสอบ พบว่า สำหรับขนาดตัวอย่างที่เปอร์เซนต์การถูกตัดทิ้ง 90% สถิติทดสอบ K ให้ค่าอำนาจการทดสอบสูงที่สุด และที่เปอร์เซนต์การถูกตัดทิ้ง 95% และ 99% สถิติทดสอบ KS, K และ CVM ให้ค่าอำนาจการทดสอบใกล้เคียงกัน

การแจกแจงคอมเพิร์ตซ์ในทุกกรณีศึกษาสามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้ ด้วยเกณฑ์ของ Cochran และเมื่อพิจารณาค่าอำนาจการทดสอบ พบว่า สำหรับทุกขนาดตัวอย่าง ที่เปอร์เซ็นต์การถูกตัดทิ้ง 90% และ 95% สถิติทดสอบ K ให้ค่าอำนาจการทดสอบสูงที่สุด และที่เปอร์เซ็นต์การถูกตัดทิ้ง 99% สถิติทดสอบ KS, K และ CVM ให้ค่าอำนาจการทดสอบใกล้เคียงกัน

Shimokawa T. และ Liao M. (1999) ได้ศึกษาเกี่ยวกับคุณสมบัติของ Kolmogorov-Smirnov (K-S), Cramervon Mises (C-M) และ Anderson-Darling (A-D) สำหรับการทดสอบ GOF สำหรับข้อมูลที่มีลักษณะค่าสุดขีด (Type-I Extreme-Value) และข้อมูลที่มีลักษณะการแจกแจงไวบูลล์มีพารามิเตอร์ 2 ตัว (2-Parameter Weibull Distribution) ประมาณค่ากลุ่มตัวอย่างโดยใช้เทคนิค GPT (Graphical Plotting Techniques) ซึ่งจะใช้วิธีการกำลังสองน้อยที่สุดประมาณค่าพารามิเตอร์ กับข้อมูลที่มีลักษณะค่าสุดขีด (Type-I Extreme-Value) และข้อมูลที่มีลักษณะการแจกแจงไวบูลล์มีพารามิเตอร์ 2 ตัว (2-Parameter Weibull Distribution) ค่าวิกฤติของ Kolmogorov-Smirnov (K-S), Cramervon Mises (C-M) และ Anderson-Darling (A-D) โดยจำลองข้อมูลจากวิธีการ Monte Carlo simulation ประกอบด้วยชุดข้อมูล 4 ชุดข้อมูล อำนาจการทดสอบของ Kolmogorov-Smirnov (K-S), Cramervon Mises (C-M) และ Anderson-Darling (A-D) เมื่อนำมาใช้กับเทคนิค GPT (Graphical Plotting Techniques) ซึ่งจากผลการทดสอบ Anderson-Darling (A-D) ให้ค่าอำนาจการทดสอบที่สูงที่สุดในทุกชุดข้อมูล

Aslan B. และ Zech G. (2002) ได้ศึกษาการเปรียบเทียบระหว่างการทดสอบไคสแควร์กับวิธีการทดสอบแบบใหม่ ได้แก่ Univariate Three-region-Test และ Multivariate Energy Test เพื่อหาการทดสอบที่เหมาะสม สำหรับทุกการแจกแจงปราศจากความเอนเอียงของกลุ่มตัวอย่าง โดยกำหนดสมมติฐานหลัก คือกลุ่มตัวอย่างมาจากฟังก์ชันความหนาแน่นของความน่าจะเป็น และสมมติฐานรองกลุ่มตัวอย่าง ไม่ได้มาจากฟังก์ชันความหนาแน่นของความน่าจะเป็นซึ่งเมื่อทำการเปรียบเทียบการทดสอบไคสแควร์กับวิธีการทดสอบแบบใหม่ มีความเอนเอียงมากน้อยเพียงใด ผลการวิจัยสรุปได้ว่าการเปรียบเทียบระหว่างการทดสอบไคสแควร์กับวิธีการทดสอบแบบใหม่ Univariate Three-region-Test เป็นวิธีที่ดีเพราะมีความเอนเอียงน้อยที่สุด

Steele, M., Chaneline, J. และ Hurst, C. (2005) ได้เสนองานวิจัยเกี่ยวกับการเปรียบเทียบอำนาจการสอบของสองประชากรโดยใช้ GOF ซึ่งจะใช้วิเคราะห์ลักษณะของข้อมูล โดยการนำข้อมูลที่จำลองขนาดตัวอย่าง 10,000 ตัว นำมาทดสอบสมมติฐานทางสถิติเกี่ยวกับพารามิเตอร์ของแต่ละการแจกแจง จะใช้วิธีการทดสอบไคสแควร์ เป็นเกณฑ์ในการตัดสินใจที่ระดับนัยสำคัญ 5% ซึ่ง โดยมีการกำหนดสมมติฐานหลักว่าข้อมูลเข้าสู่การแจกแจง Uniform มากกว่า 10 cell และกำหนดสมมติฐานรองเข้าสู่การแจกแจง Increasing และ Triangular v ผลการวิจัยสรุปได้ว่า ปฏิเสธ H_0 : ที่ระดับนัยสำคัญ 0.05

ดังนั้นแสดงข้อมูลส่วนใหญ่จะมีการแจกแจงสู่เข้าสู่ Increasing และ Triangular v ที่ระดับนัยสำคัญ 0.05

Chernobai A., Rachev S. และ Fabozzi F.(2005) ได้ศึกษาการเปรียบเทียบวิธีการทดสอบ GOF โดยประยุกต์ในกรณีที่มีการสูญหายของข้อมูลซึ่งเป็นข้อมูลการประกันภัย เพื่อหาความถูกต้องในการวิเคราะห์สมมติฐานของข้อมูล โดยใช้วิธีการทดสอบ KS, Kuiper, Anderson-Darling และ Cramer-Von Mises โดยใช้เทคนิคในการตรวจสอบ GOF ของจำนวนการแจกแจงของกลุ่มตัวอย่าง โดยใช้เกณฑ์การตัดสินใจเมื่อคำนวณความเสี่ยงของข้อมูลที่มีความสมบูรณ์ กับความเสี่ยงของข้อมูลที่สูญหายบางส่วน ว่าวิธีไหนมีความคลาดเคลื่อนจากความเป็นจริงน้อยที่สุดแสดงว่าวิธีการทดสอบนั้นดีที่สุด ผลการทดสอบสรุปว่าวิธีการทดสอบ Anderson-Darling เป็นวิธีการทดสอบที่มีความคลาดเคลื่อนจากความเป็นจริงน้อยที่สุดเมื่อนำไปเทียบกับข้อมูลที่มีความสมบูรณ์

บทที่ 3

วิธีการดำเนินการวิจัย

ในการวิจัยครั้งนี้เป็นการวิจัยเชิงทดลองเพื่อทำการทดสอบภาวะสารูปดี (Goodness of Fit Test : GOF) สำหรับการจำแนกการแจกแบบลอการมัล (Log-normal Distribution) และการแจกแบบไวบูลล์ (Weibull Distribution) ซึ่งสถิติทดสอบที่ใช้ในการวิจัย ได้แก่ Anderson Darling Test Statistic, Kolmogorov -Smirnov Test Statistic, Kuiper Test Statistic และ Ratio of Maximum Likelihoods Test Statistic โดยสถานการณ์ต่างๆ ที่ใช้ในการวิจัยนั้นสร้างโดยใช้โปรแกรม R และแสดงความสามารถในการควบคุมค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 และการเปรียบเทียบค่าอำนาจการทดสอบของตัวสถิติทดสอบที่กล่าวในข้างต้น

3.1 แผนการทดลอง

กำหนดสถานการณ์ต่าง ๆ สำหรับศึกษาเปรียบเทียบความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 และอำนาจการทดสอบของตัวสถิติทดสอบ ทั้ง 4 ตัว ดังนี้

3.1.1 เลือกสุ่มตัวอย่างจากประชากรโดยกำหนดประชากรให้มีการแจกแจงดังต่อไปนี้

- การแจกแจงแบบลอการมัล (Log-normal Distribution) โดยสนใจศึกษาเมื่อพารามิเตอร์ ดังนี้

$$\mu=1 \text{ และ } \sigma=0.5$$

$$\mu=1 \text{ และ } \sigma=1$$

$$\mu=1 \text{ และ } \sigma=1.5$$

$$\mu=1 \text{ และ } \sigma=2$$

$$\mu=1 \text{ และ } \sigma=3$$

- การแจกแจงแบบไวบูลล์ (Weibull Distribution) โดยสนใจศึกษาเมื่อพารามิเตอร์ ดังนี้

$$\beta=1 \text{ และ } \alpha=0.5$$

$$\beta=1 \text{ และ } \alpha=1$$

$$\beta=1 \text{ และ } \alpha=1.5$$

$$\beta=1 \text{ และ } \alpha=2$$

$$\beta=1 \text{ และ } \alpha=3$$

3.1.2 การกำหนดขนาดของตัวอย่าง (sample size) กำหนดให้ขนาดตัวอย่างที่ใช้ในการศึกษามี 4 ระดับ คือ 20, 30, 40 และ 50

3.2 ขั้นตอนในการวิจัย

ขั้นตอนในการวิจัยแบ่งออกเป็น 4 ขั้นตอน ดังนี้

3.2.1 สร้างการแจกแจงของประชากรตามลักษณะที่กำหนดในแผนการทดลอง

3.2.2 คำนวณค่าตัวสถิติทดสอบ 4 ตัว

3.2.3 คำนวณค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1

3.2.4 คำนวณค่าอำนาจการทดสอบ

รายละเอียดแต่ละขั้นตอนมีดังนี้

3.2.1 สร้างการแจกแจงของประชากรตามลักษณะที่กำหนดในแผนการทดลอง

สร้างลักษณะการแจกแจงของประชากรตามลักษณะที่กำหนดในแผนการทดลองใช้โปรแกรม R รายละเอียดในการสร้างการแจกแจงแบบต่าง ๆ เป็นดังนี้

3.2.1.1 การผลิตเลขสุ่มที่มีการแจกแจงแบบไวบูลล์

การแจกแจงแบบไวบูลล์มีฟังก์ชันความหนาแน่นอยู่ในรูปของ

$$g(x;\beta,\alpha) = \frac{\alpha}{\beta} \left(\frac{x}{\beta}\right)^{\alpha-1} \exp\left[-\left(\frac{x}{\beta}\right)^{\alpha}\right], \quad x > 0, \beta > 0, \alpha > 0$$

เมื่อ β เป็น scale parameter

α เป็น shape parameter

ค่าคาดหวังและความแปรปรวนของการแจกแจงแบบไวบูลล์ คือ $\beta\Gamma(1+1/\alpha)$ และ $\beta^2[\Gamma(1+2/\alpha) - \Gamma^2(1+1/\alpha)]$ ตามลำดับ

การสร้างตัวแปรสุ่มให้มีการแจกแจงไวบูลล์ โดยใช้โปรแกรม R ซึ่งคำสั่งในการสร้างตัวแปรให้มีการแจกแจงไวบูลล์ จำนวน n ตัว ที่พารามิเตอร์ $\beta = b$ และ $\alpha = c$ คือ `rweibull(n, c, b)` ดังตัวอย่างต่อไปนี้

```
> rweibull(20, 0.5, 1)
```

```
[1] 5.1755583635 0.0366120499 0.2183311546 0.0772656155 0.0001667598
```

```
[6] 0.0654551852 0.0045679611 0.0071867150 0.0378979935 0.0225855857
```

```
[11] 0.5057724261 0.3996091438 0.8400628308 1.5302615961 0.1395647976
```

```
[16] 0.0159737345 8.1828089643 0.2581195054 2.9580998477 1.5356178105
```

3.2.1.2 การผลิตเลขสุ่มที่มีการแจกแจงแบบลอกนอร์มัล

การแจกแจงแบบลอกนอร์มัลมีฟังก์ชันความหนาแน่นอยู่ในรูปของ

$$f(x; \mu, \sigma) = \frac{1}{x\sqrt{2\pi\sigma^2}} \exp\left[-(\ln x - \mu)^2 / 2\sigma^2\right], \quad x > 0, \sigma > 0$$

เมื่อ σ เป็น scale parameter

μ เป็น shape parameter

ค่าคาดหวังและความแปรปรวนของการแจกแจงแบบลอกนอร์มัลคือ $\exp(2\mu + \sigma^2)$ และ $\exp(2\mu + \sigma^2) * [\exp(\sigma^2) - 1]$ ตามลำดับ

การสร้างตัวแปรสุ่มให้มีการแจกแจงลอกนอร์มัล โดยใช้โปรแกรม R ซึ่งคำสั่งในการสร้างตัวแปรให้มีการแจกแจงไวบูลล์จำนวน n ตัว ที่พารามิเตอร์ $\sigma=b$ และ $\mu=c$ คือ `rlnorm(n, c, b)` ดังตัวอย่างต่อไปนี้

```
> rlnorm(20,1,0.5)
[1] 1.4042463 2.2420105 3.1535787 3.6548522 4.8233718 1.9393909 5.2188760
[8] 8.0466736 2.5247879 2.7275269 3.5761696 1.2750229 1.0618689 1.8664705
[15] 0.9356022 1.1947580 4.6557036 2.8146571 3.9825933 8.5046730
```

3.2.2 คำนวณค่าตัวสถิติทดสอบ 4 ตัว

เมื่อสร้างตัวแปรสุ่มตามลักษณะการแจกแจง ขนาดตัวอย่างตามแผนการทดลองที่กำหนด แล้วนำข้อมูลที่ได้อมาคำนวณค่าสถิติทดสอบแต่ละวิธีดังนี้

3.2.2.1 สถิติทดสอบ Anderson-Darling

$$A = \left[- \sum_{i=1}^n \frac{(2i-1)}{n} \{ \ln(F_0(x_i)) + \ln(1 - F_0(x_{n+1-i})) \} \right] - n$$

3.2.2.2 สถิติทดสอบ Kolmogorov - Smirnov

$$KS = \max(D^+, D^-)$$

$$\text{เมื่อ } D^+ = \max_{i=1, \dots, n} \left\{ \frac{i}{n} - F_0(x_i) \right\}$$

$$\text{และ } D^- = \max_{i=1, \dots, n} \left\{ F_0(x_i) - \frac{i-1}{n} \right\}$$

3.2.2.3 สถิติทดสอบ Kuiper

$$V = D_N^+ + D_N^-$$

$$\text{เมื่อ } D^+ = \max_{i=1, \dots, n} \left\{ \frac{i}{n} - F_0(x_i) \right\}$$

$$\text{และ } D^- = \max_{i=1, \dots, n} \left\{ F_0(x_i) - \frac{i-1}{n} \right\}$$

3.2.2.4 สถิติทดสอบ Ratio of Maximum Likelihoods

ก) ตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงไวบูลล์

$$(RML)^{\frac{1}{n}} = (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_{i=1}^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}}$$

เมื่อ $\hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n}$ และ $\hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n}$

ข) ตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงลอกนอร์มัล

$$(RML)^{\frac{1}{n}} = (2\pi e \hat{\sigma}^2)^{\frac{1}{2}} \left[\prod_{i=1}^n x_i f_w(x_i; \hat{b}, \hat{c}) \right]^{\frac{1}{n}}$$

เมื่อ $\hat{\sigma}^2 = \frac{\sum_{i=1}^n (\ln x_i - \hat{\mu})^2}{n}$ และ $\hat{\mu} = \frac{\sum_{i=1}^n (\ln x_i)}{n}$

เมื่อได้ค่าของตัวสถิติทดสอบแต่ละตัวให้ในการหาค่าวิกฤตของสถิติทดสอบนั้น ๆ

3.2.3 คำนวณค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1

การคำนวณค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 สามารถหาได้โดยการนับจำนวนครั้งของการปฏิเสธสมมติฐานว่าง แล้วหารด้วยจำนวนรอบทั้งหมด ซึ่งมีขั้นตอนดังนี้

3.2.3.1 ตั้งสมมติฐานของการทดลอง

ก) ถ้าสมมติฐานว่าง (H_0) คือ การแจกแจงลอกนอร์มัล

สมมติฐานแย้ง (H_1) คือ การแจกแจงไวบูลล์

ข) ถ้าสมมติฐานว่าง (H_0) คือ การแจกแจงไวบูลล์

สมมติฐานแย้ง (H_1) คือ การแจกแจงลอกนอร์มัล

3.2.3.2 กำหนดลักษณะการแจกแจงของประชากร ในการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 การแจกแจงของประชากรจะเป็นไปตามสมมติฐานว่าง (H_0) ที่ตั้งไว้ และพารามิเตอร์ที่สนใจศึกษาในแต่ละการแจกแจง

3.2.3.3 กำหนดขนาดตัวอย่าง 4 ขนาด คือ 20, 30, 40 และ 50

3.2.3.4 สุ่มตัวอย่างจากที่กำหนดในข้อ 2.33.2 ถึง 3.2.3.3

3.2.3.5 คำนวณค่าสถิติทดสอบที่ใช้ทดสอบการแจกแจงตามสมมติฐานว่าง (H_0)

3.2.3.6 นำค่าสถิติทดสอบที่คำนวณได้จากข้อ 3.2.3.5 มาเปรียบเทียบกับค่าวิกฤตในแต่ละตัวสถิติทดสอบ ที่ระดับนัยสำคัญ (α_1) 0.1, 0.05, 0.01 แล้วนับจำนวนครั้งที่ปฏิเสธสมมติฐานว่าง (H_0)

3.2.3.7 ทำซ้ำข้อ 3.2.3.4 ถึง 3.2.3.6 จำนวน 1,000 รอบ

3.2.3.8 คำนวณค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยนับจำนวนครั้งที่ปฏิเสธว่างทั้งหมด แล้วนำมารวมด้วยจำนวนรอบทั้งหมด 1,000 รอบ

เมื่อคำนวณค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้นำค่าที่ได้มาเปรียบเทียบกับค่าที่ได้ในแต่ละสถานการณ์กับเกณฑ์ของ Cochran และ Bradley แล้วพิจารณาสถานการณ์ใดสามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้ และมีสถานการณ์ใดที่ไม่สามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้

3.2.4 คำนวณค่าอำนาจการทดสอบ

การคำนวณค่าอำนาจการทดสอบคล้ายกับการคำนวณค่าความคลาดเคลื่อนประเภทที่ 1 คือนับจำนวนครั้งที่การปฏิเสธสมมติฐานว่าง แล้วหารด้วยจำนวนรอบทั้งหมด แต่ลักษณะการแจกแจงของประชากรเป็นไปตามสมมติฐานแย้ง (H_1) ซึ่งมีขั้นตอนการทำงานดังนี้

3.2.4.1 ตั้งสมมติฐานของการทดลอง

ก) ถ้าสมมติฐานว่าง (H_0) คือ การแจกแจงลอการิธึม

สมมติฐานแย้ง (H_1) คือ การแจกแจงไวบูลล์

ข) ถ้าสมมติฐานว่าง (H_0) คือ การแจกแจงไวบูลล์

สมมติฐานแย้ง (H_1) คือ การแจกแจงลอการิธึม

3.2.4.2 กำหนดลักษณะการแจกแจงของประชากร ในการหาค่าอำนาจการทดสอบ การแจกแจงของประชากรจะเป็นไปตามสมมติฐานแย้ง (H_1) ที่ตั้งไว้ และพารามิเตอร์ที่สนใจศึกษาในแต่ละการแจกแจง

3.2.4.3 กำหนดขนาดตัวอย่าง 4 ขนาด คือ 20, 30, 40 และ 50

3.2.4.4 สุ่มตัวอย่างจากที่กำหนดในข้อ 3.2.4.2 ถึง 3.2.4.3

3.2.4.5 คำนวณค่าสถิติทดสอบที่ใช้ทดสอบการแจกแจงตามสมมติฐานว่าง (H_0)

3.2.4.6 นำค่าสถิติทดสอบที่คำนวณได้จากข้อ 3.2.4.5 มาเปรียบเทียบกับค่าวิกฤตในแต่ละตัวสถิติทดสอบ ที่ระดับนัยสำคัญ (α_1) 0.1, 0.05, 0.01 แล้วนับจำนวนครั้งที่ปฏิเสธสมมติฐานว่าง (H_0)

3.2.4.7 ทำซ้ำข้อ 3.2.4.4 ถึง 3.2.4.6 จำนวน 1,000 รอบ

3.2.4.8 คำนวณค่าอำนาจการทดสอบ โดยนับจำนวนครั้งที่ปฏิเสธสมมติฐานว่าง (H_0) ทั้งหมด แล้วนำมารวมด้วยจำนวนรอบทั้งหมด 1,000 รอบ

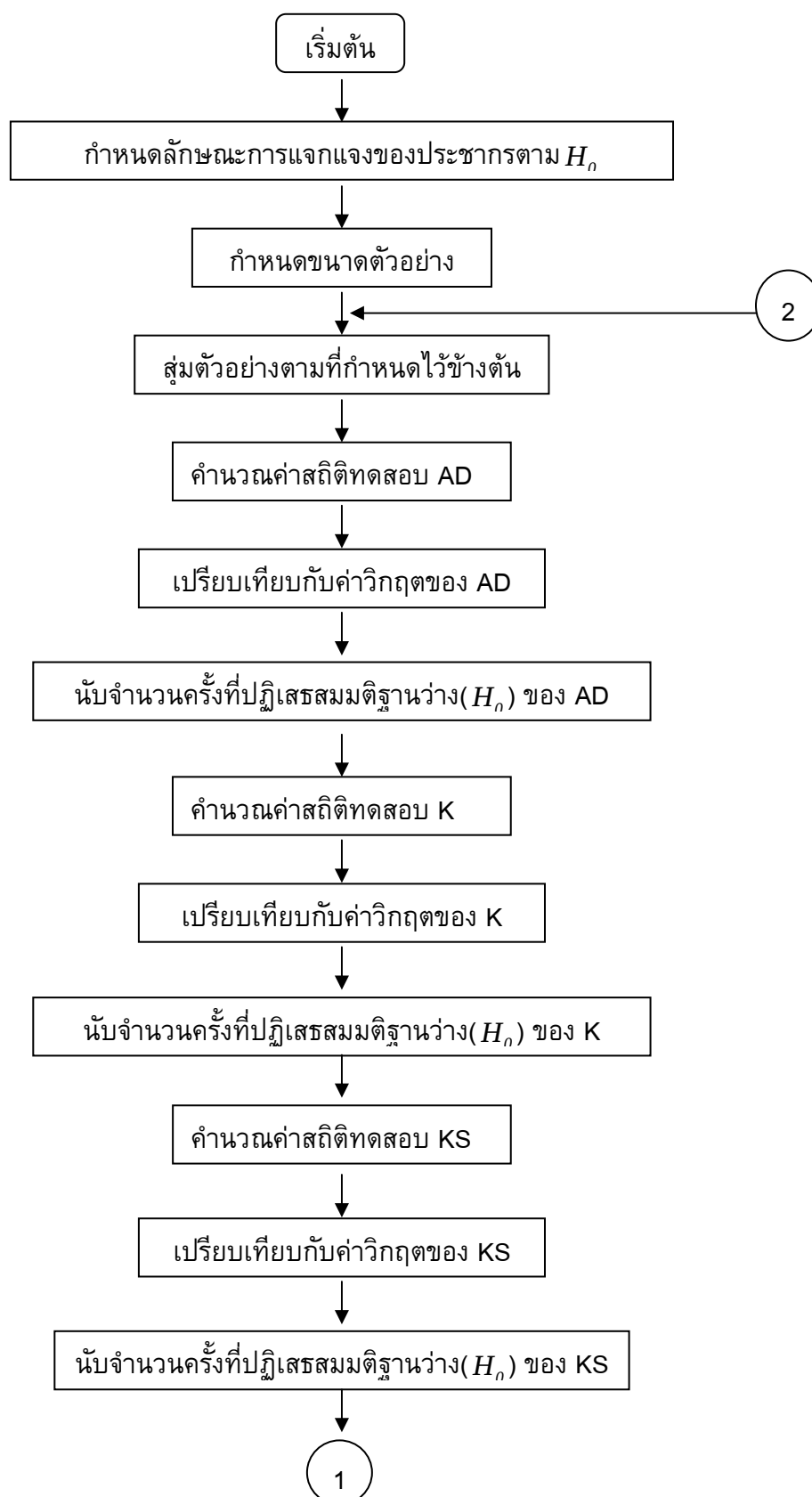
เมื่อคำนวณค่าอำนาจการทดสอบได้แล้ว จะเปรียบเทียบกับค่าที่ได้ของแต่ละค่าสถิติทดสอบในแต่ละกรณี

3.3 ขั้นตอนการทำงานของโปรแกรม

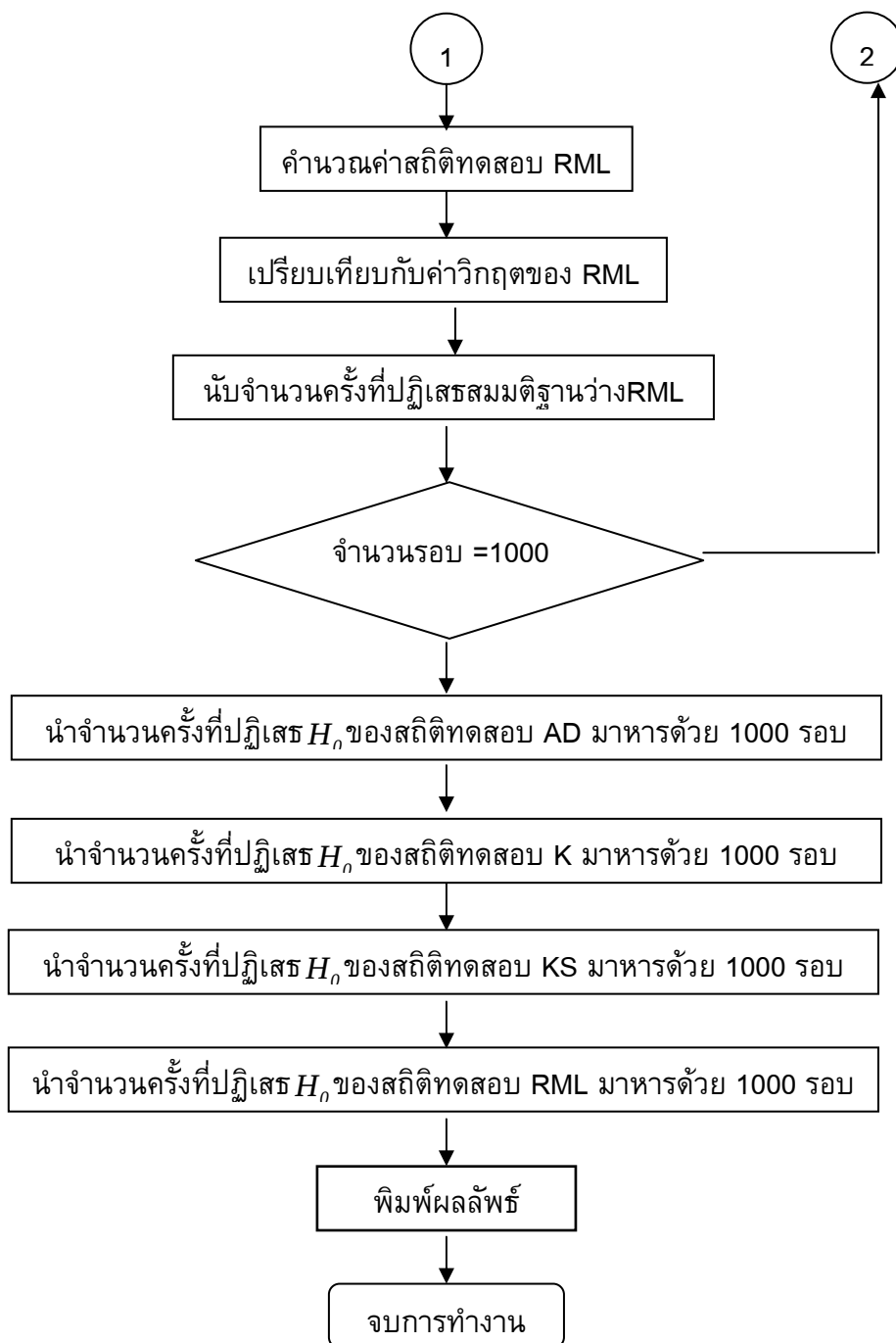
สามารถเขียนผังงานแสดงขั้นตอนในการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ตามรูปที่ 3.1 และผังงานแสดงขั้นตอนการหาค่าอำนาจการทดสอบ ได้ตามรูปที่ 3.2 ดังรายละเอียดต่อไปนี้

ภาพที่ 3-1 แผนผังแสดงขั้นตอนการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1

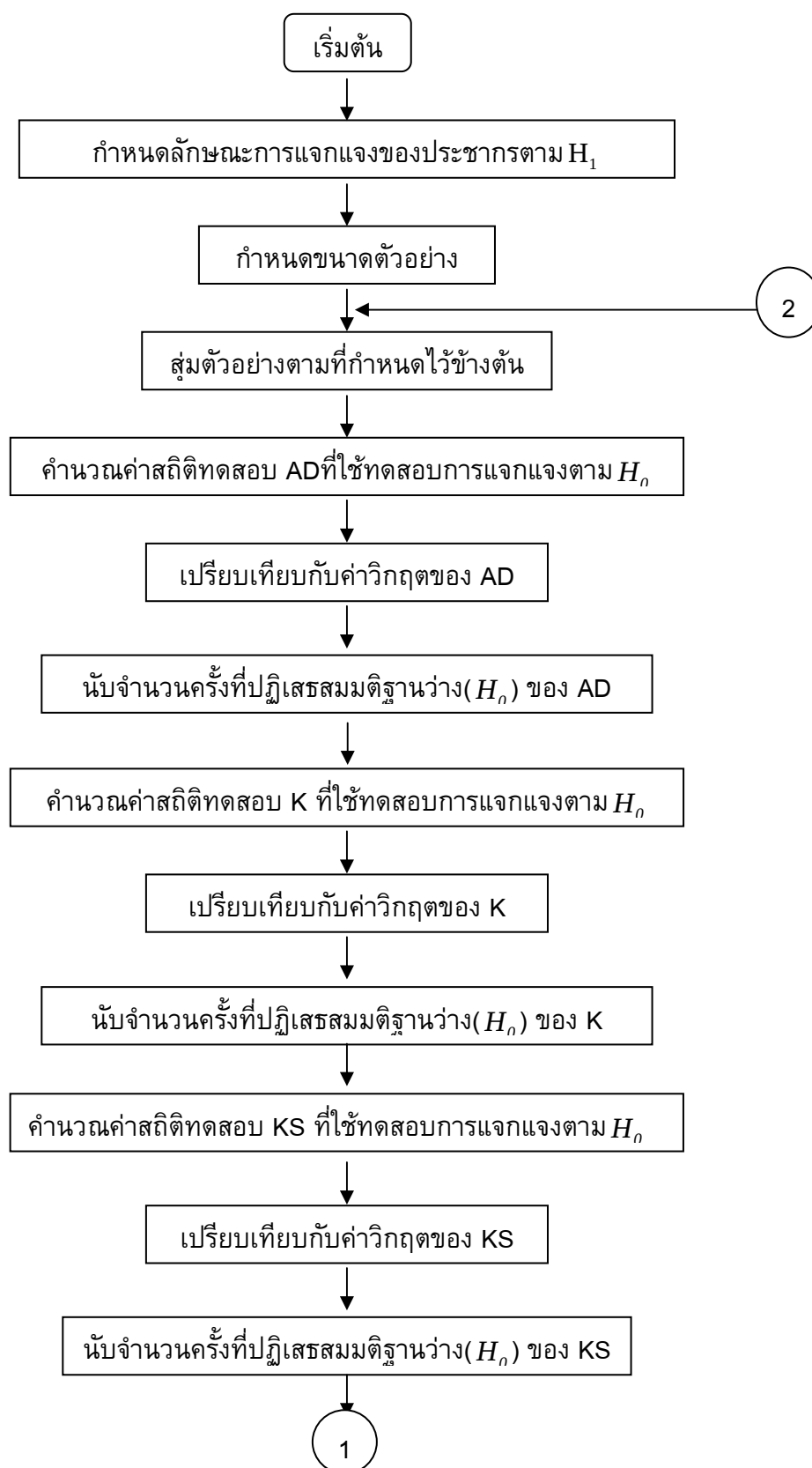
ภาพที่ 3-2 แผนผังแสดงขั้นตอนการหาค่าอำนาจการทดสอบ



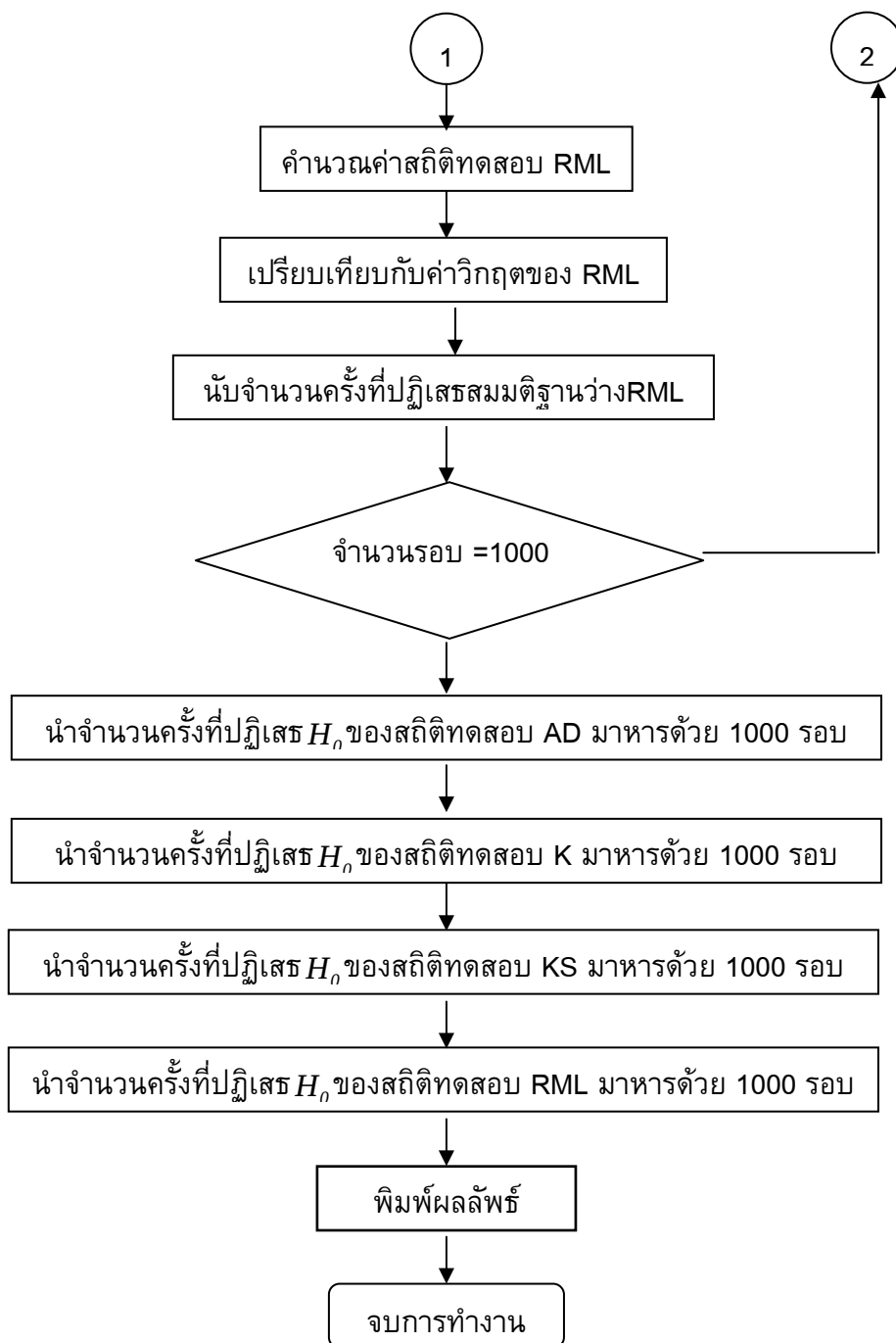
ภาพที่ 3-1 แผนผังแสดงขั้นตอนการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1



ภาพที่ 3-1 (ต่อ)



ภาพที่3-2 แผนผังแสดงขั้นตอนการหาค่าอำนาจการทดสอบ



ภาพที่3-2 (ต่อ)

บทที่ 4

ผลการวิจัย

การวิจัยครั้งนี้ ผู้วิจัยมีความสนใจศึกษาวิธีการทดสอบภาวะสารูปดี (Goodness of Fit Test) ในการจำแนกการแจกแบบลอการิธึมและการแจกแจงแบบไวบูลล์ โดยจะศึกษา Anderson Darling Test Statistic, Kuiper Test Statistic, Ratio of Maximum Likelihoods Test Statistic และ Kolmogorov -Smirnov Test Statistic โดยพิจารณาจากความสามารถในการควบคุมค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 แล้วทำการเปรียบเทียบประสิทธิภาพของวิธีการทดสอบ ด้วยค่าอำนาจของการทดสอบของสถิติทดสอบทั้ง 4 วิธีข้างต้น

ในการทดสอบภาวะสารูปดี (Goodness of Fit Test) จะต้องมีการตั้งสมมติฐานขึ้นแล้วนำมาทำการทดสอบสมมติฐาน โดยทั่วไปพบว่าผลการทดสอบที่ได้อาจเกิดความคลาดเคลื่อน 2 ประเภท คือ ความคลาดเคลื่อนประเภทที่ 1 (Type I error) และความคลาดเคลื่อนประเภทที่ 2 (Type II error) ดังแสดงได้ในตารางต่อไปนี้

ตารางที่ 4-1 แสดงความคลาดเคลื่อนในการทดสอบสมมติฐานทางสถิติ

สมมติฐานว่า	ผลการทดลอง	
	ปฏิเสธ H_0	ยอมรับ H_0
H_0		
จริง	ความคลาดเคลื่อนประเภทที่ 1 (α_1)	ตัดสินใจถูก ($1 - \alpha_1$)
เท็จ	ค่าอำนาจการทดสอบ ($1 - \beta_1$)	ความคลาดเคลื่อนประเภทที่ 2 (β_1)

ในการทดสอบสมมติฐานทางสถิติ ต้องการให้ความน่าจะเป็นที่จะเกิดความคลาดเคลื่อนประเภทที่ 1 (α_1) และความน่าจะเป็นที่จะเกิดความคลาดเคลื่อนประเภทที่ 2 (β_1) มีค่าน้อยที่สุด เพื่อให้ค่าอำนาจการทดสอบ ($1 - \beta_1$) มีค่ามากที่สุด แต่ถ้าวัด α_1 จะทำให้ β_1 เพิ่มขึ้น และถ้าวัด β_1 จะทำให้ α_1 เพิ่มขึ้น ดังนั้นในการเปรียบเทียบอำนาจการทดสอบที่จะควบคุม α_1 โดยพิจารณาความสามารถ ในการควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 แล้วจึงเปรียบเทียบอำนาจการทดสอบ ของตัวสถิติทดสอบที่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้เท่านั้น

ในการวิจัยครั้งนี้ได้เสนอผลการวิจัยเป็น 2 ขั้นตอน คือ

ขั้นตอนที่ 1 เปรียบเทียบความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากการทดลอง โดยใช้ตัวสถิติทดสอบทั้ง 4 วิธีข้างต้น เพื่อสรุปว่าตัวสถิติทดสอบใดสามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อกำหนดระดับนัยสำคัญ (α_1) ของการทดสอบเท่ากับ 0.01, 0.05 และ 0.1

สถิติทดสอบที่มีความแกร่ง (Robustness) จะต้องไม่แสดงความไว (sensitive) ต่อการทดสอบในกรณีที่ข้อมูลไม่เป็นไปตามข้อตกลงเบื้องต้นของการทดสอบ ซึ่งจะพิจารณาได้จากค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 (τ) ที่ได้จากการทดลองเปรียบเทียบกับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 (α_1) ที่กำหนดไว้ โดยใช้เกณฑ์ในการพิจารณาความสามารถในการควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของ Cochran (1954 : อ้างโดย ศิริรัตน์ 2539, 35) และเกณฑ์ของ Bradley (1978 : อ้างโดย ศิริรัตน์ 2539, 35) ซึ่งในการวิจัยครั้งนี้จะพิจารณาควบคู่กันไปด้วยรายละเอียดสำหรับแต่ละเกณฑ์ดังนี้

1) เกณฑ์ของ Cochran

เกณฑ์นี้ใช้ตัดสินใจความสามารถในการควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยกำหนดให้ τ คือ ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ที่เกิดจากการทดลองสถิติทดสอบจะควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ก็ต่อเมื่อ

ค่า τ อยู่ในช่วง $[0.007, 0.015]$ ที่ระดับนัยสำคัญ 0.01

ค่า τ อยู่ในช่วง $[0.04, 0.06]$ ที่ระดับนัยสำคัญ 0.05

ค่า τ อยู่ในช่วง $[0.08, 0.12]$ ที่ระดับนัยสำคัญ 0.1

2) เกณฑ์ของ Bradley

เกณฑ์นี้ใช้ตัดสินใจความสามารถในการควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยกำหนดให้ τ คือ ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ที่เกิดจากการทดลองสถิติทดสอบจะควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ก็ต่อเมื่อค่า τ อยู่ในช่วง $[0.05\alpha_1, 1.5\alpha_1]$ ที่ระดับนัยสำคัญ α_1 นั่นคือ

ค่า τ อยู่ในช่วง $[0.005, 0.015]$ ที่ระดับนัยสำคัญ 0.01

ค่า τ อยู่ในช่วง $[0.025, 0.075]$ ที่ระดับนัยสำคัญ 0.05

ค่า τ อยู่ในช่วง $[0.05, 0.15]$ ที่ระดับนัยสำคัญ 0.1

จากผลการทดลอง ถ้าความน่าจะเป็นของความคลาดเคลื่อนที่ 1 ของการทดลองอยู่นอกเหนือช่วงที่กำหนด หรือนอกเหนือช่วงที่กำหนด หรือการทดลองนั้นไม่สามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้ สามารถแบ่งได้เป็น 2 กรณี คือ

1. กรณีที่ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 สูงกว่าขอบเขตบนที่กำหนด จะถือว่า การทดสอบนั้นมีความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 มากกว่าระดับนัยสำคัญที่กำหนด ($\tau > \alpha_1$)

2. กรณีที่ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ต่ำกว่าขอบเขตบนที่กำหนด จะถือว่า การทดสอบนั้นมีความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 น้อยกว่าระดับนัยสำคัญที่กำหนด ($\tau < \alpha_1$)

ในกรณีที่ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของการทดลองอยู่ในขอบเขตที่ระบุสำหรับแต่ละเกณฑ์ที่กำหนด จะถือว่า การทดสอบนั้น มีความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 เท่ากับค่า α_1 ที่กำหนด และสามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้

สำหรับการนำเสนอความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ซึ่งได้จากการทดลองจะนำเสนอในรูปของตารางและกราฟ ซึ่งจะแบ่งการนำเสนอออกเป็น 2 การแจกแจงด้วยกัน คือการแจกแจงไวบูลล์ 2 พารามิเตอร์ และการแจกแจงลอกนอร์มัล

ขั้นตอนที่ 2_ เปรียบเทียบอำนาจการทดสอบ ซึ่งเป็นสิ่งที่ใช้พิจารณาในการเลือกตัวสถิติทดสอบ คือเลือกใช้ตัวสถิติทดสอบที่ปฏิเสธสมมุติฐาน H_0 มากที่สุด เมื่อสมมุติฐาน H_0 นั้นผิด ซึ่งหมายความว่า ตัวสถิติทดสอบนั้นให้อำนาจการทดสอบสูงสุด ในการวิจัยครั้งนี้จะเปรียบเทียบอำนาจการทดสอบของตัวสถิติทดสอบทั้ง 4 วิธี เมื่อประชากรมีการแจกแจงไวบูลล์ 2 พารามิเตอร์ และการแจกแจงลอกนอร์มัล ซึ่งจะนำเสนอด้วยตารางและรูปกราฟ

4.1 การเปรียบเทียบความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1

ผลการวิจัยสำหรับการแจกแจงไวบูลล์

การนำเสนอความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากการทดลองของสถิติทดสอบทั้ง 4 ตัว สำหรับการแจกแจงไวบูลล์ ได้แสดงเป็นตารางที่ 4-2 – 4-6

ตารางที่ 4-2 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$

n	α_1	AD	K	KS	RML
20	0.1	0.094	0.068*	0.073*	0.106
	0.05	0.048	0.030*	0.038*	0.050
	0.01	0.007	0.005*	0.006*	0.009
30	0.1	0.107	0.081	0.090	0.102
	0.05	0.055	0.041	0.049	0.053
	0.01	0.018**	0.015	0.013	0.013
40	0.1	0.101	0.104	0.110	0.096
	0.05	0.054	0.045	0.060	0.047
	0.01	0.009	0.012	0.012	0.011
50	0.1	0.122*	0.106	0.103	0.129*
	0.05	0.065*	0.053	0.047	0.061*
	0.01	0.018**	0.013	0.015	0.017**

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-2 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

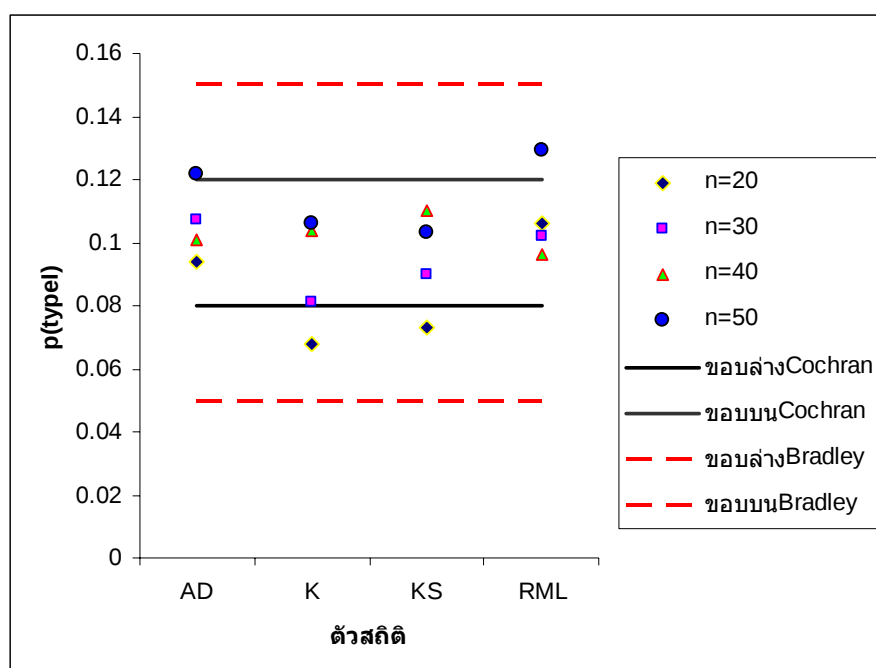
1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01 และ กรณีที่ $n = 50$ ระดับนัยสำคัญ 0.05 และ 0.01
- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ทุกระดับนัยสำคัญ ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ทุกระดับนัยสำคัญ
- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 50$ ทุกระดับนัยสำคัญ

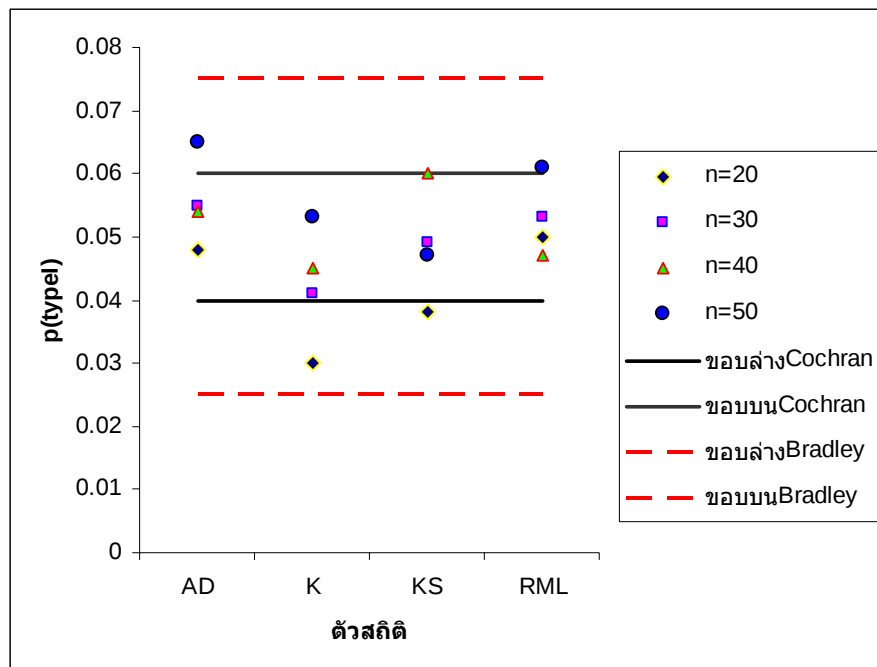
2) เกณฑ์ของ Bradley

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01 และ กรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01
- ตัวสถิติทดสอบ K สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

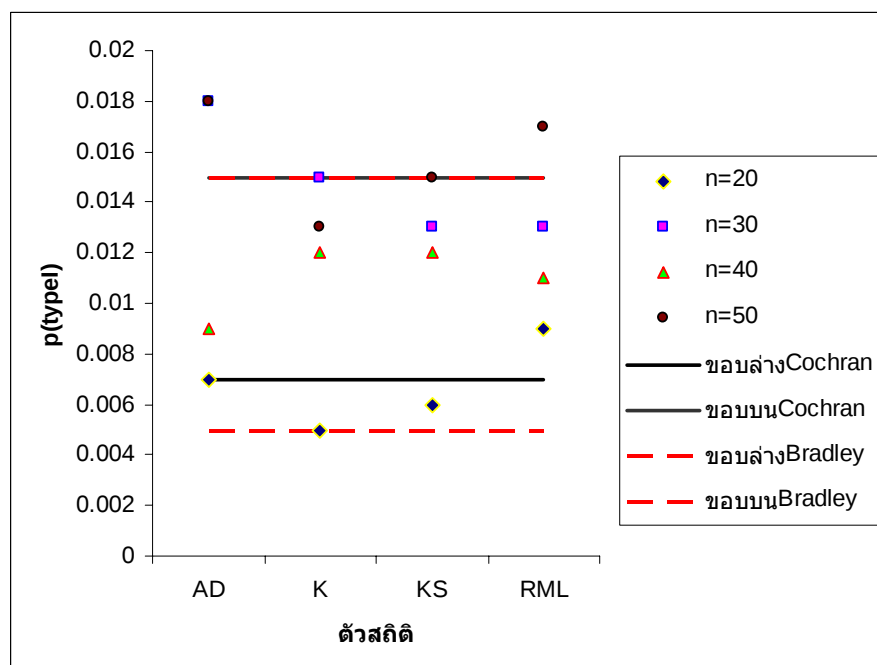
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-1 - 4-3 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-1 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-2 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-3 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-3 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1$

n	α_1	AD	K	KS	RML
20	0.1	0.101	0.075*	0.072*	0.115
	0.05	0.057	0.036*	0.040	0.060
	0.01	0.017**	0.009	0.009	0.012
30	0.1	0.112	0.079*	0.098	0.107
	0.05	0.053	0.037*	0.038 *	0.065*
	0.01	0.010	0.004**	0.004**	0.009
40	0.1	0.097	0.092	0.100	0.100
	0.05	0.041	0.044	0.041	0.048
	0.01	0.007	0.009	0.011	0.014
50	0.1	0.093	0.080	0.082	0.098
	0.05	0.043	0.034*	0.042	0.052
	0.01	0.008	0.011	0.009	0.014

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-3 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ 0.05 ในกรณีที่ $n = 30$ ทุกระดับนัยสำคัญ และกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.05

- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.05 และ 0.01

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.05

2) เกณฑ์ของ Bradley

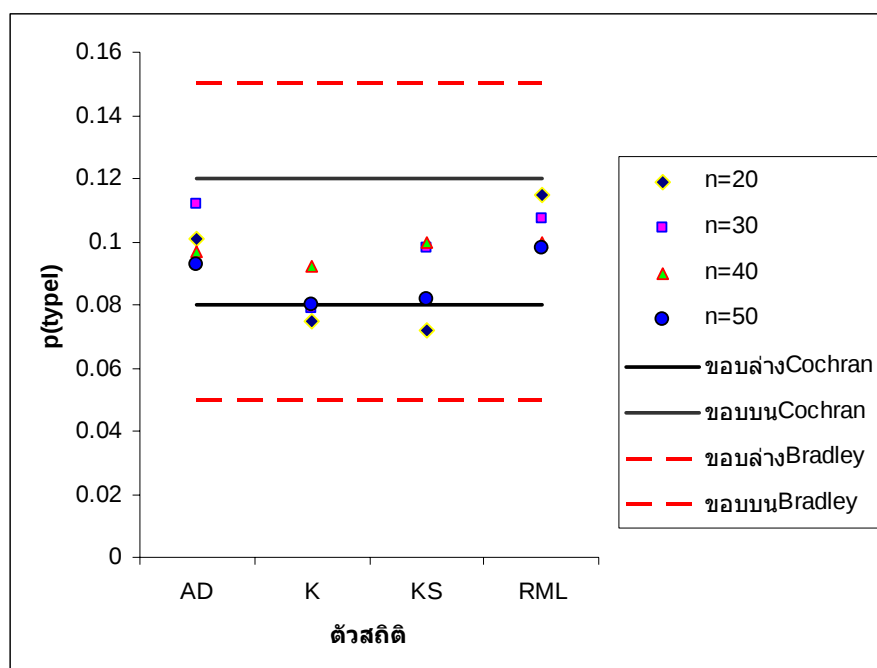
- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01

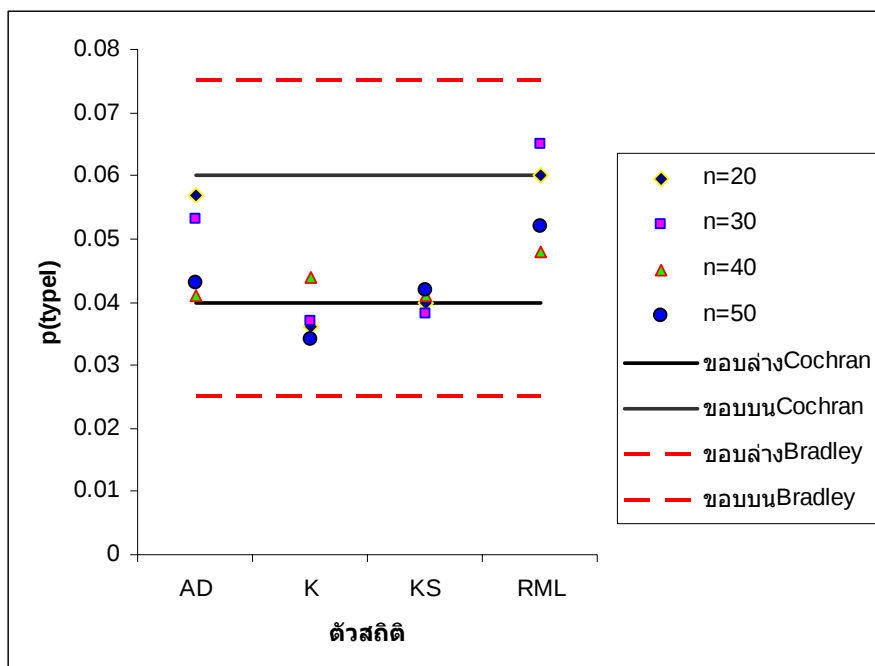
- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

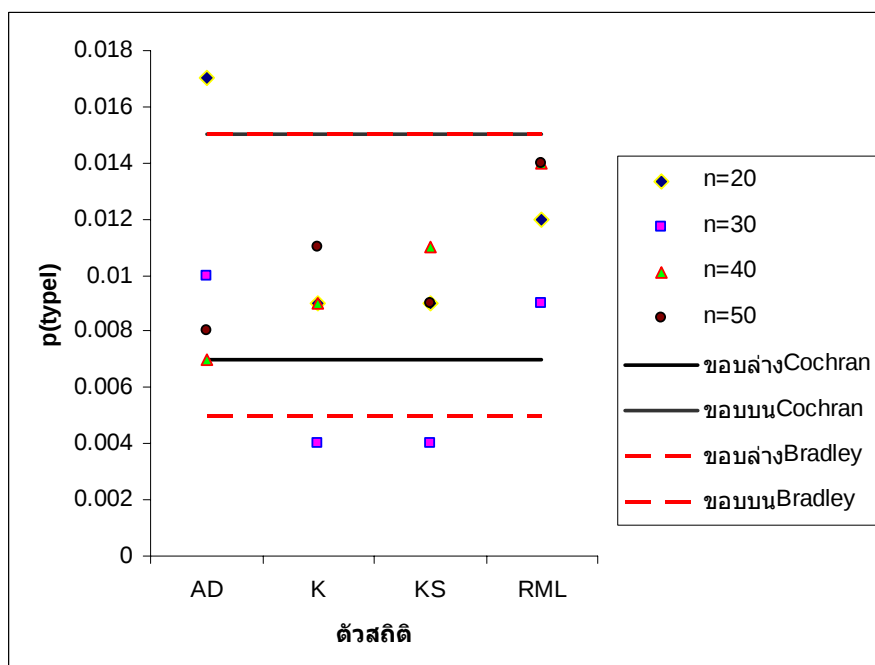
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-4 – 4-6 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-4 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-5 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-6 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-4 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1.5$

n	α_1	AD	K	KS	RML
20	0.1	0.108	0.075*	0.071*	0.092
	0.05	0.050	0.038*	0.038*	0.043
	0.01	0.009	0.008	0.008	0.009
30	0.1	0.117	0.083	0.089	0.111
	0.05	0.052	0.051	0.049	0.058
	0.01	0.013	0.010	0.0012**	0.011
40	0.1	0.098	0.083	0.095	0.100
	0.05	0.047	0.040	0.037*	0.049
	0.01	0.008	0.003**	0.006*	0.015
50	0.1	0.094	0.095	0.099	0.110
	0.05	0.055	0.050	0.048	0.046
	0.01	0.014	0.011	0.013	0.014

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-4 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ 0.05 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ 0.05 ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01 และ ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.05 และ 0.01

- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

2) เกณฑ์ของ Bradley

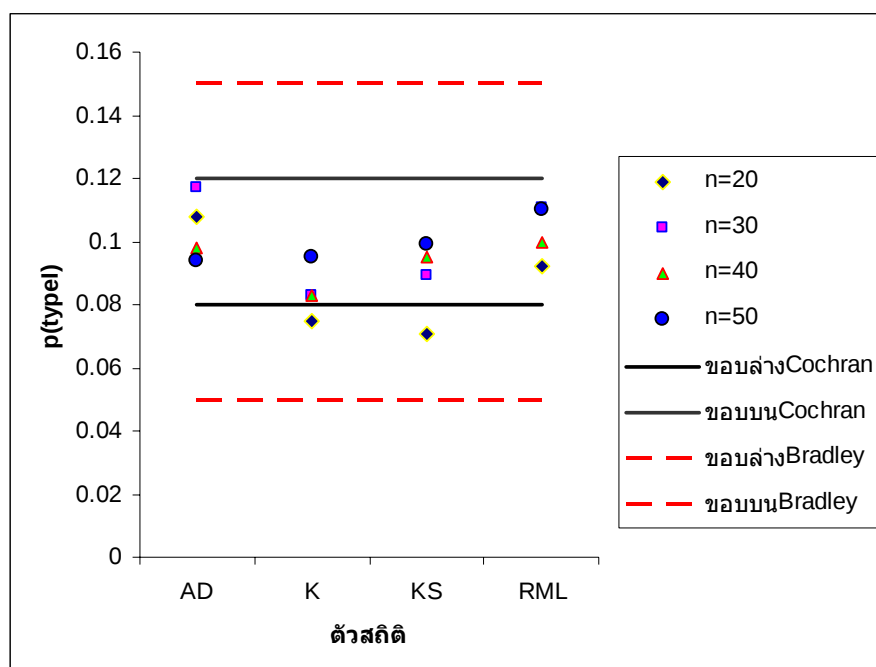
- ตัวสถิติทดสอบ AD สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01

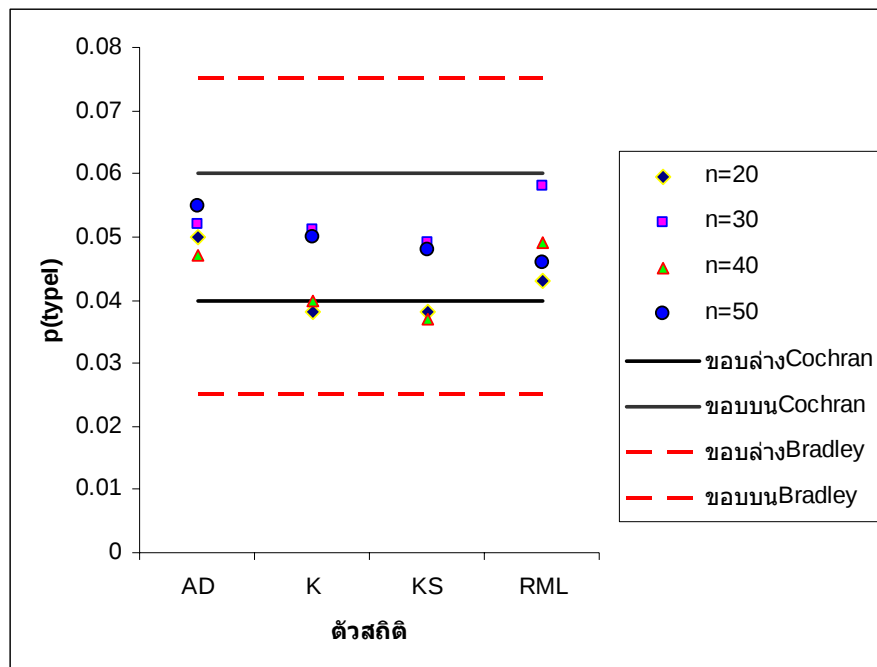
- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

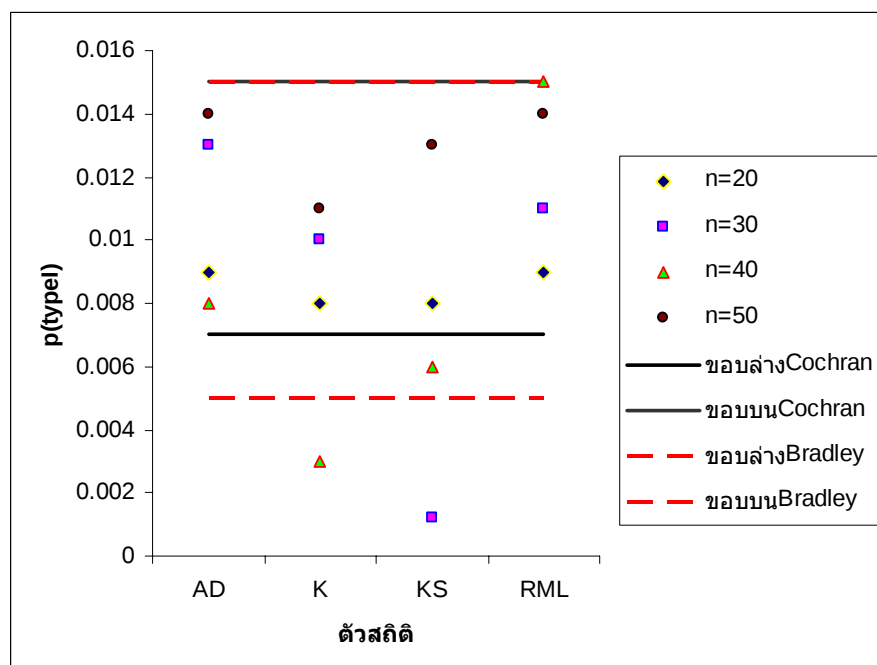
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-7 – 4-9 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-7 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-8 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-9 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1.5$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-5 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha = 2$

n	α_1	AD	K	KS	RML
20	0.1	0.104	0.076*	0.083	0.095
	0.05	0.052	0.035*	0.041	0.045
	0.01	0.013	0.011	0.010	0.013
30	0.1	0.104	0.083	0.083	0.094
	0.05	0.054	0.032*	0.040	0.043
	0.01	0.011	0.006*	0.006*	0.013
40	0.1	0.102	0.069*	0.071*	0.122*
	0.05	0.048	0.027*	0.029*	0.067*
	0.01	0.008	0.007	0.005*	0.016**
50	0.1	0.109	0.096	0.088	0.106
	0.05	0.059	0.053	0.042	0.044
	0.01	0.010	0.008	0.009	0.017**

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-5 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ 0.05 ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.05 และ 0.01 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.1 และ 0.05

- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01 และในกรณีที่ $n = 40$ ทุกระดับนัยสำคัญ

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ในกรณีที่ $n = 40$ ทุกระดับนัยสำคัญ และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

2) เกณฑ์ของ Bradley

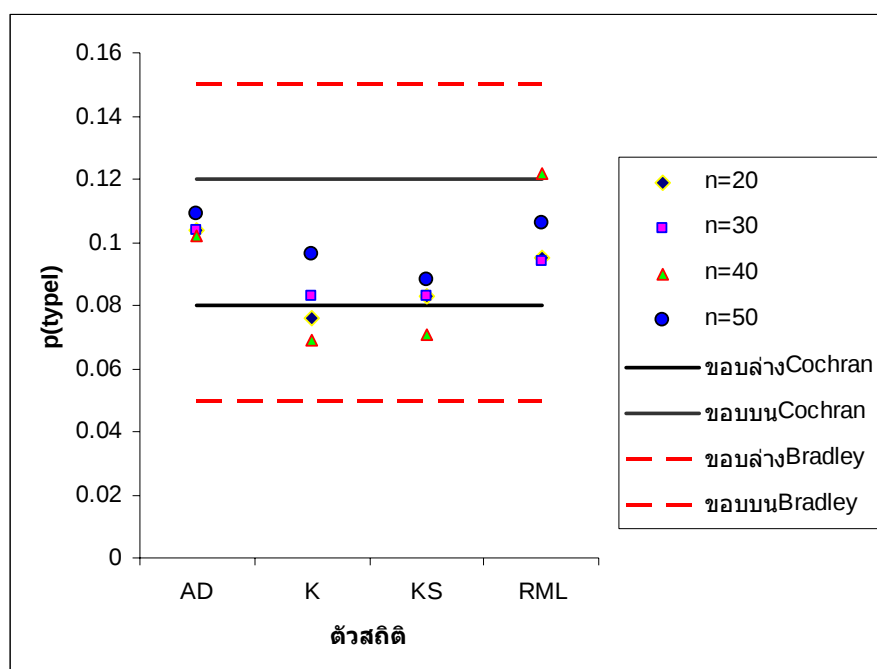
- ตัวสถิติทดสอบ AD สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ K สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

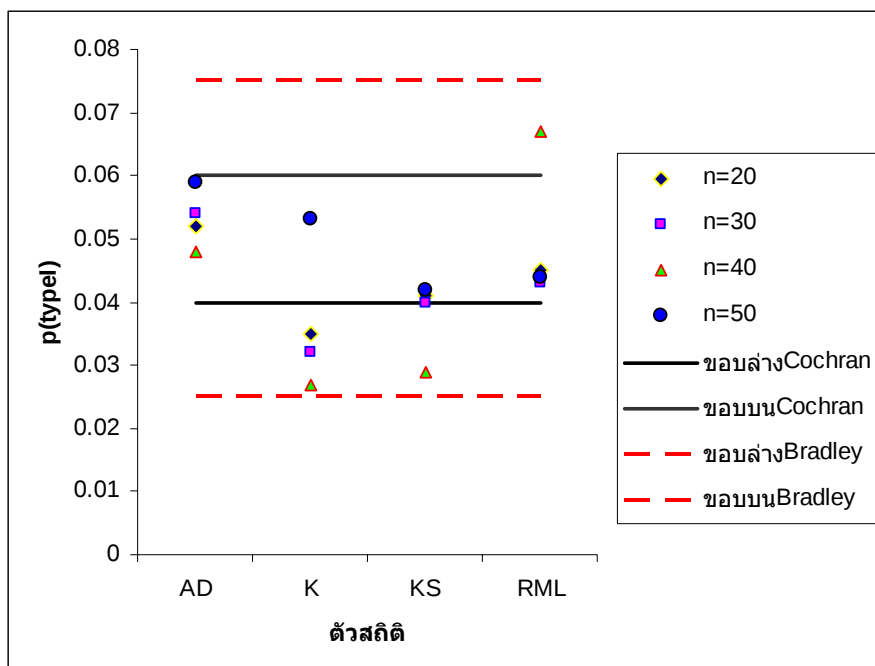
- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

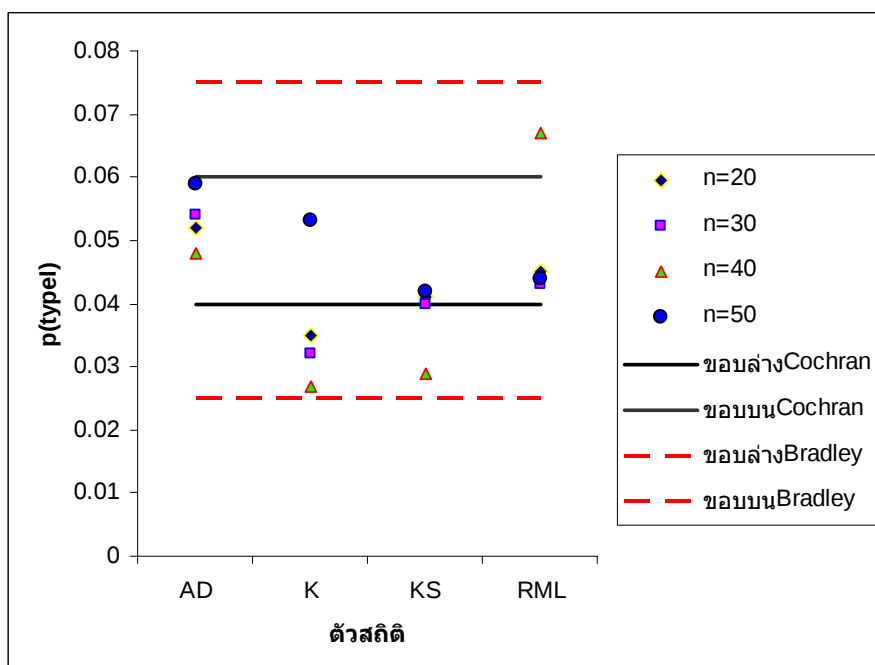
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-10 – 4-11 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-10 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลส์ $\beta = 1$ และ $\alpha = 2$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-11 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-12 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-6 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ

$$\alpha = 3$$

n	α_1	AD	K	KS	RML
20	0.1	0.119	0.077*	0.085	0.113
	0.05	0.067*	0.041	0.043	0.061*
	0.01	0.016**	0.007	0.008	0.014
30	0.1	0.093	0.077*	0.082	0.108
	0.05	0.044	0.031*	0.040	0.055
	0.01	0.013	0.007	0.011	0.009
40	0.1	0.108	0.084	0.086	0.105
	0.05	0.053	0.036*	0.050	0.059
	0.01	0.007	0.004**	0.004**	0.013
50	0.1	0.102	0.084	0.096	0.100
	0.05	0.052	0.045	0.046	0.049
	0.01	0.012	0.011	0.012	0.014

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-6 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.05 และ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.1 และ 0.05 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.05 และ 0.01

- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.05

2) เกณฑ์ของ Bradley

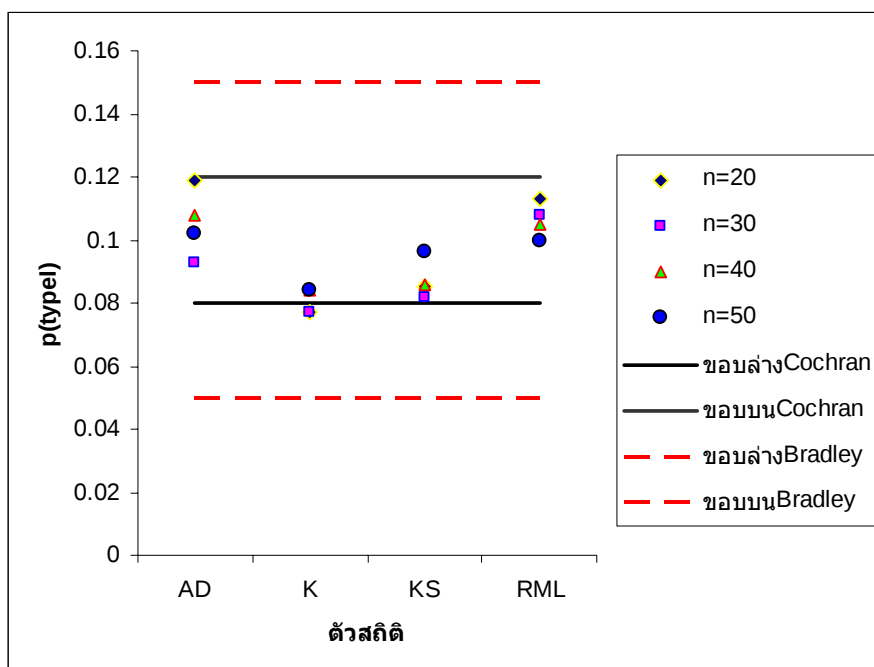
- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01

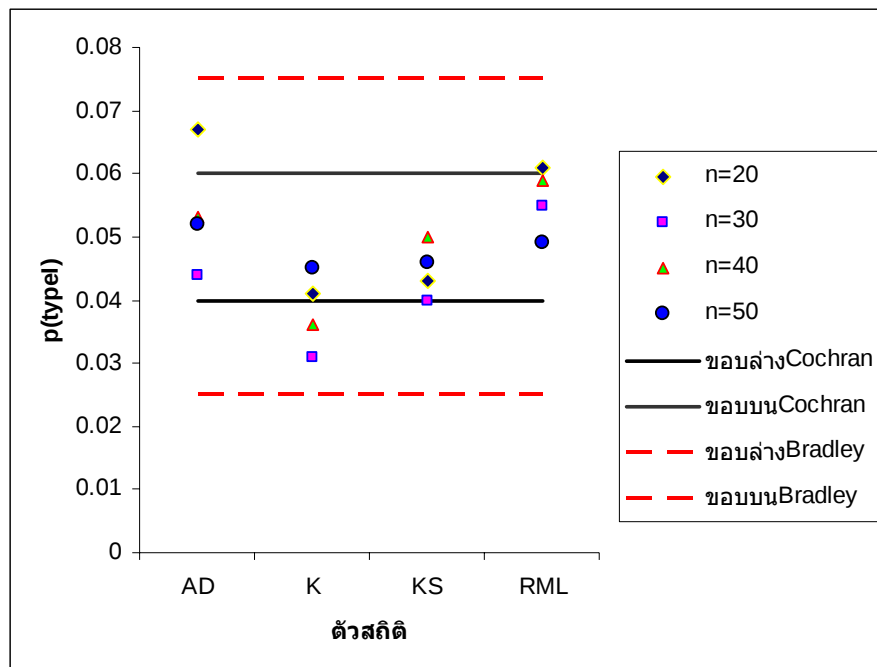
- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

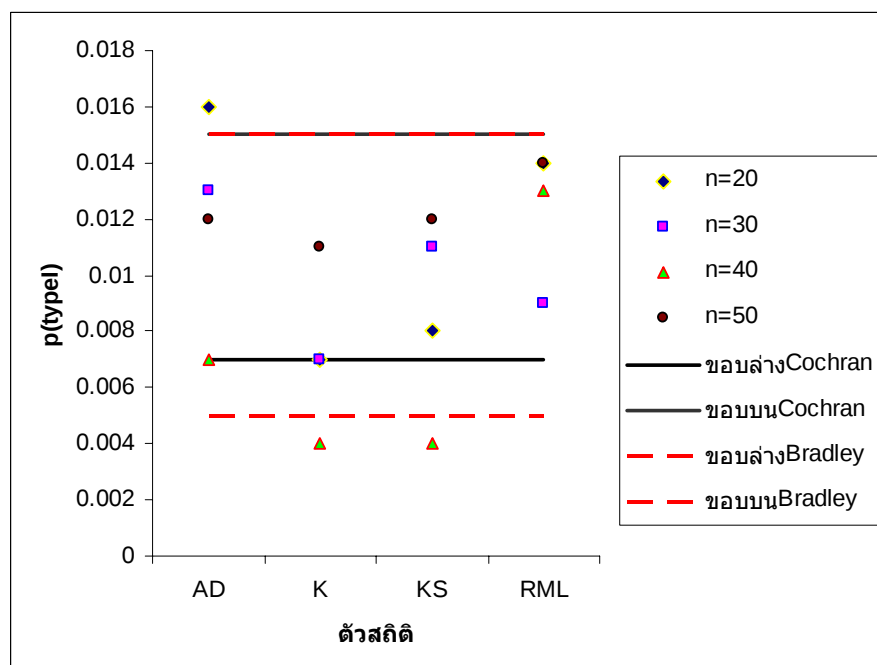
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-13 - 4-15 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-13 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลส์ $\beta = 1$ และ $\alpha = 3$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-14 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-15 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$ ที่ $\alpha_1 = 0.01$

ผลการวิจัยสำหรับการแจกแจงลอการิธึม

การนำเสนอความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากการทดลองของสถิติทดสอบทั้ง 4 ตัว สำหรับการแจกแจงลอการิธึม ได้แสดงเป็นตารางที่ 4-7 – 4-11

ตารางที่ 4-7 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 0.5$

n	α_1	AD	K	KS	RML
20	0.1	0.131*	0.115	0.108	0.118
	0.05	0.067*	0.063*	0.049	0.040
	0.01	0.012	0.015	0.012	0.015
30	0.1	0.106	0.106	0.101	0.108
	0.05	0.053	0.055	0.046	0.054
	0.01	0.014	0.011	0.014	0.007
40	0.1	0.101	0.095	0.096	0.096
	0.05	0.049	0.050	0.047	0.047
	0.01	0.010	0.009	0.006*	0.008
50	0.1	0.091	0.093	0.090	0.109
	0.05	0.047	0.048	0.045	0.053
	0.01	0.016**	0.015	0.012	0.011

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-7 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

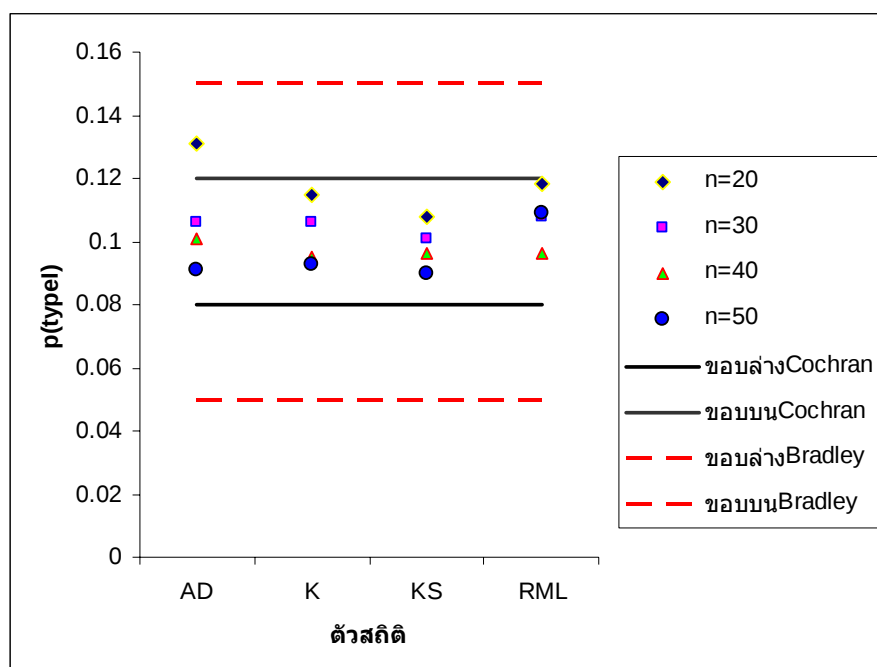
- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ 0.05 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.05
- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01
- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

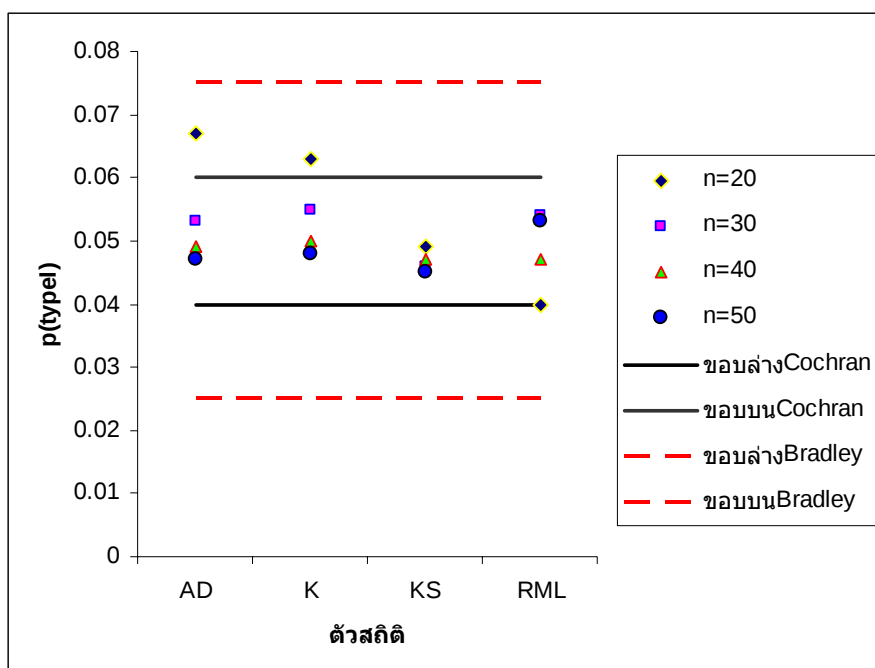
2) เกณฑ์ของ Bradley

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01
- ตัวสถิติทดสอบ K สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

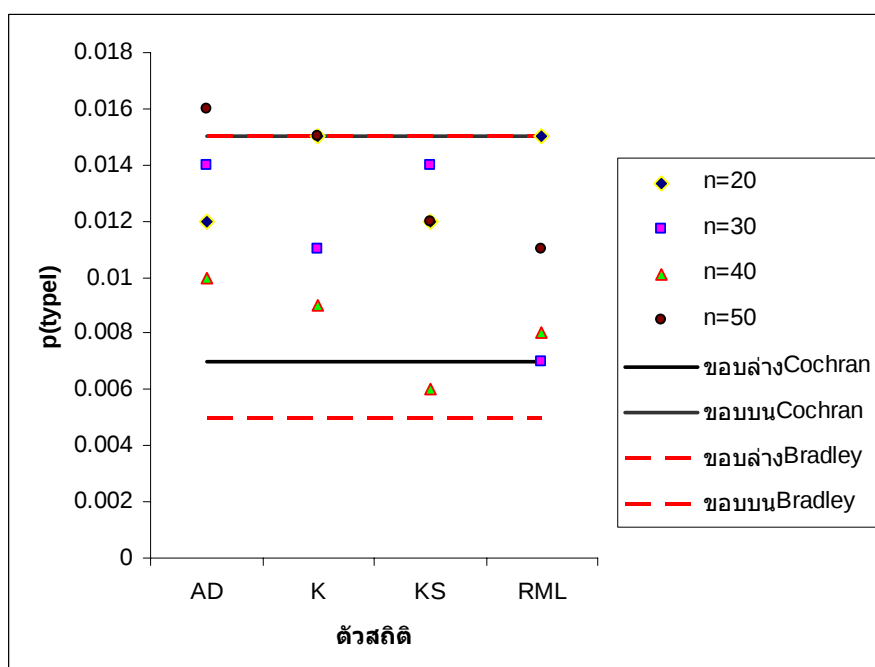
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-16 – 4-18 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-16 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-17 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-18 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\sigma = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-8 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$

n	α_1	AD	K	KS	RML
20	0.1	0.128*	0.100	0.097	0.116
	0.05	0.061*	0.052	0.045	0.045
	0.01	0.009	0.006*	0.010	0.018**
30	0.1	0.105	0.097	0.092	0.100
	0.05	0.064*	0.055	0.047	0.048
	0.01	0.011	0.007	0.009	0.013
40	0.1	0.098	0.086	0.090	0.108
	0.05	0.046	0.036*	0.041	0.054
	0.01	0.005*	0.006*	0.007	0.008
50	0.1	0.104	0.102	0.096	0.105
	0.05	0.059	0.058	0.051	0.061*
	0.01	0.008	0.016**	0.011	0.006*

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-8 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.1 และ 0.05 ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.05 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01 ในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.05 และ 0.01 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.05 และ 0.01

2) เกณฑ์ของ Bradley

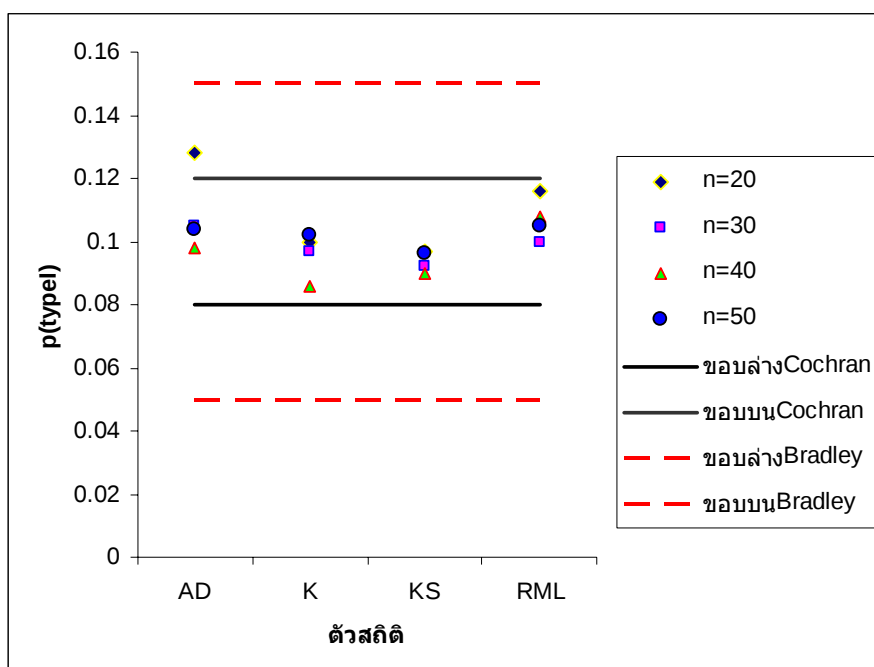
- ตัวสถิติทดสอบ AD สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

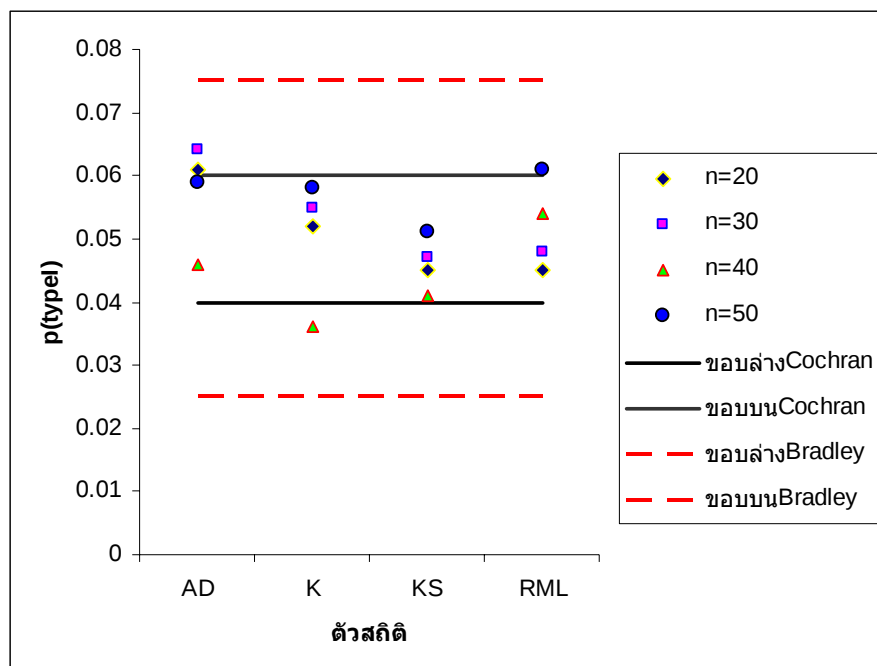
ซึ่งจะแสดงผลการเปรียบเทียบ ค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-19 – 4-21 จำแนกตามระดับนัยสำคัญ



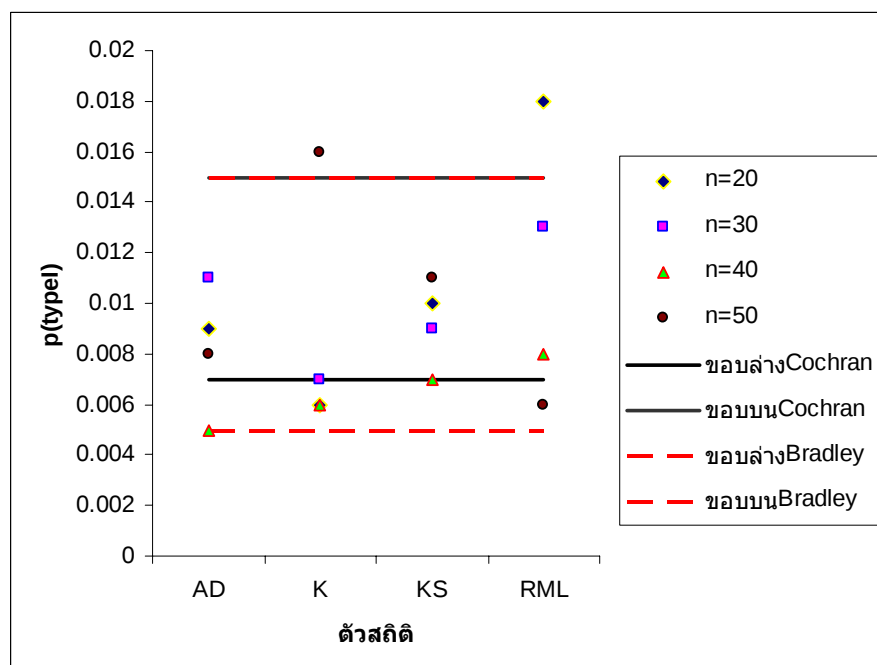
ภาพที่ 4-19 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ

AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$

ที่ $\alpha_1 = 0.1$



ภาพที่ 4-20 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-21 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-9 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1.5$

n	α_1	AD	K	KS	RML
20	0.1	0.103	0.096	0.092	0.110
	0.05	0.043	0.045	0.037*	0.030*
	0.01	0.004**	0.007	0.005*	0.006*
30	0.1	0.114	0.097	0.095	0.107
	0.05	0.051	0.047	0.044	0.042
	0.01	0.010	0.010	0.008	0.009
40	0.1	0.115	0.097	0.090	0.110
	0.05	0.063*	0.041	0.044	0.056
	0.01	0.011	0.009	0.010	0.011
50	0.1	0.105	0.104	0.081	0.095
	0.05	0.042	0.053	0.041	0.057
	0.01	0.009	0.010	0.009	0.006*

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-9 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

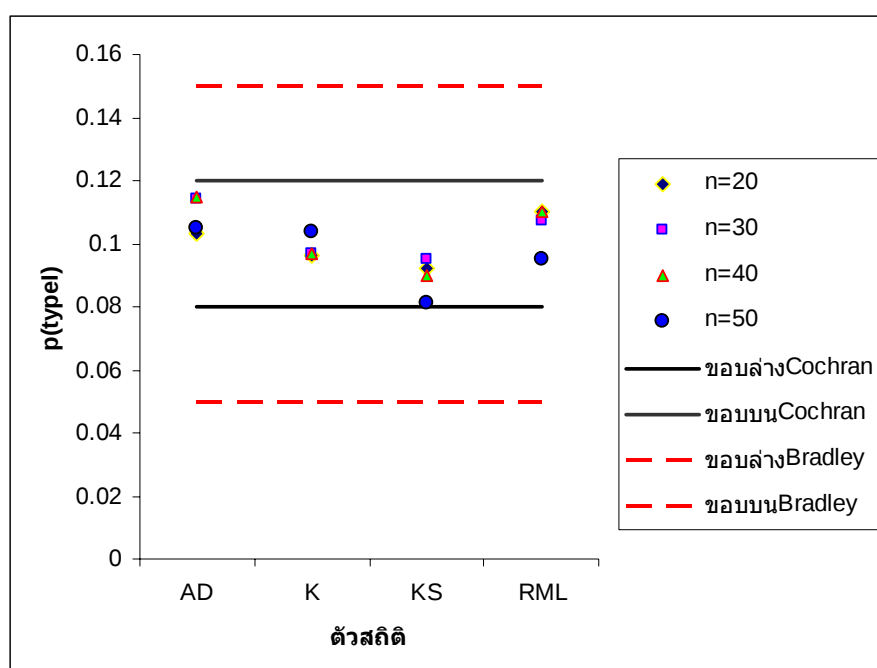
- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.05
- ตัวสถิติทดสอบ K สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.05 และ 0.01

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.05 และ 0.01 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

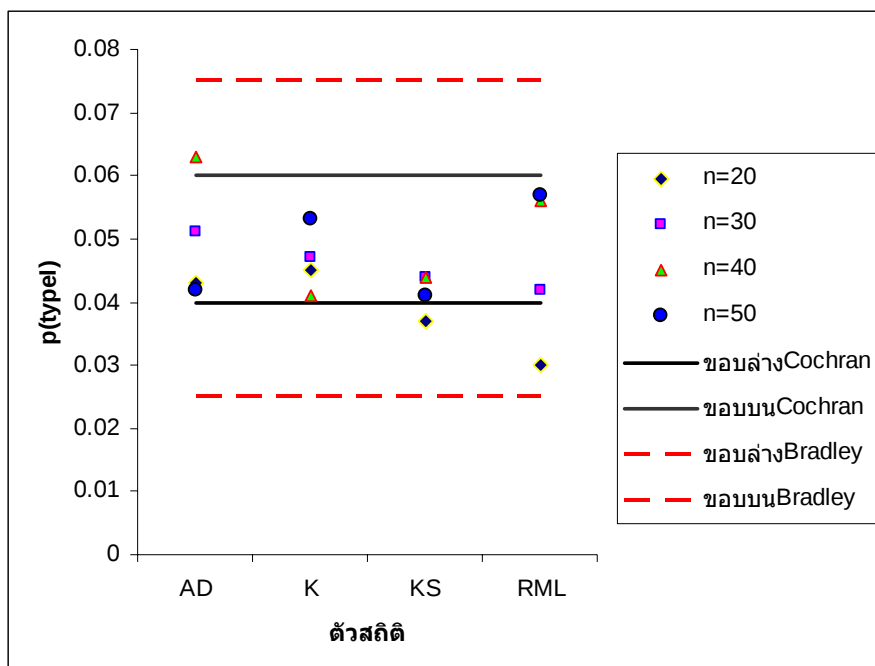
2) เกณฑ์ของ Bradley

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01
- ตัวสถิติทดสอบ K สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี
- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

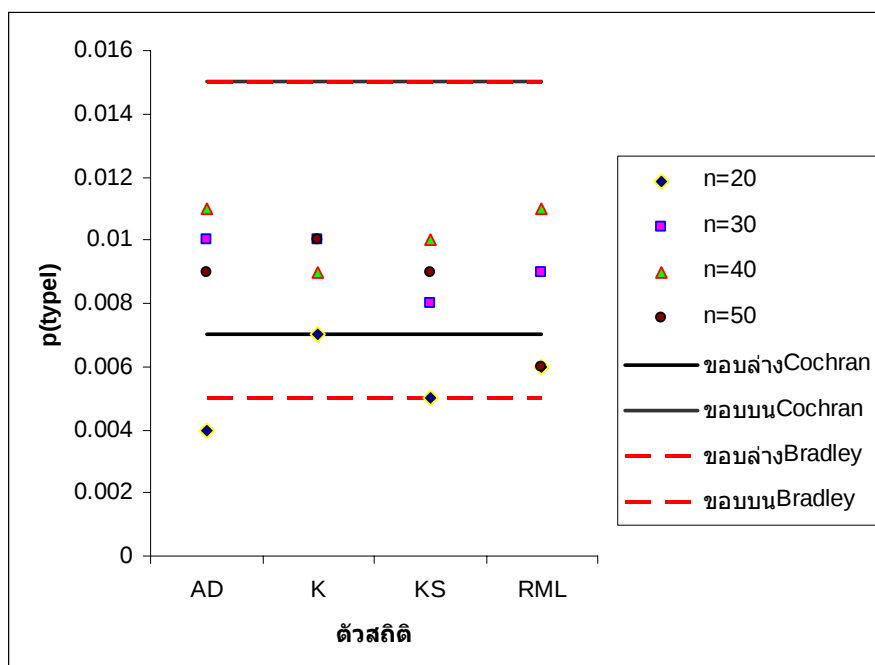
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-22- 4-24 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-22 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-23 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-24 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-10 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$

n	α_1	AD	K	KS	RML
20	0.1	0.100	0.098	0.093	0.092
	0.05	0.047	0.055	0.047	0.034
	0.01	0.014	0.011	0.014	0.009
30	0.1	0.115	0.110	0.094	0.102
	0.05	0.063*	0.059	0.056	0.056
	0.01	0.017**	0.017**	0.013	0.007
40	0.1	0.095	0.075 *	0.092	0.110
	0.05	0.045	0.035*	0.040	0.051
	0.01	0.013	0.008	0.010	0.012
50	0.1	0.119	0.104	0.110	0.132*
	0.05	0.064*	0.051	0.050	0.077**
	0.01	0.010	0.008	0.008	0.015

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4-10 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01 และ 0.05 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.05

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.1 และ 0.05

- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.1 และ 0.05

2) เกณฑ์ของ Bradley

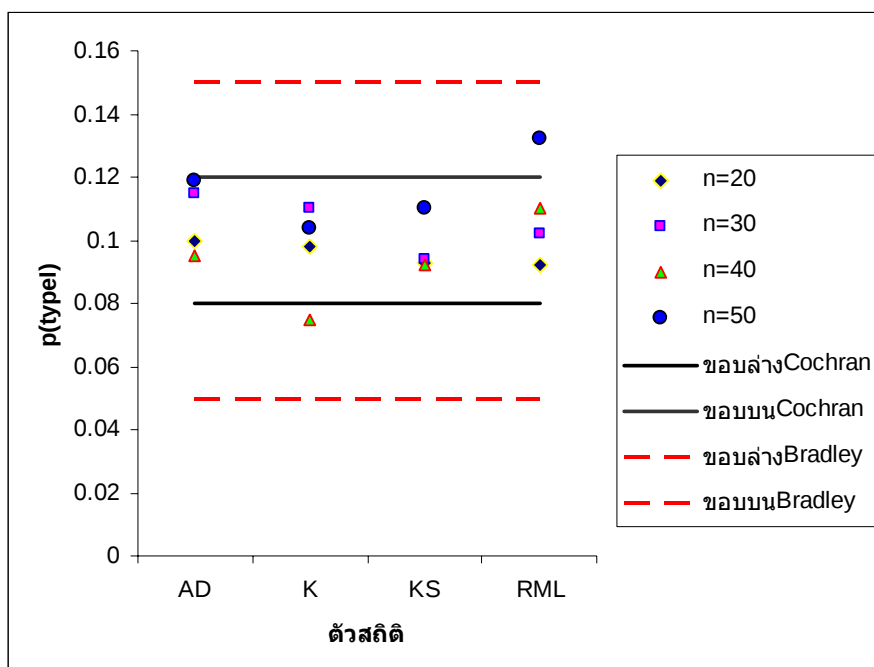
- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.05

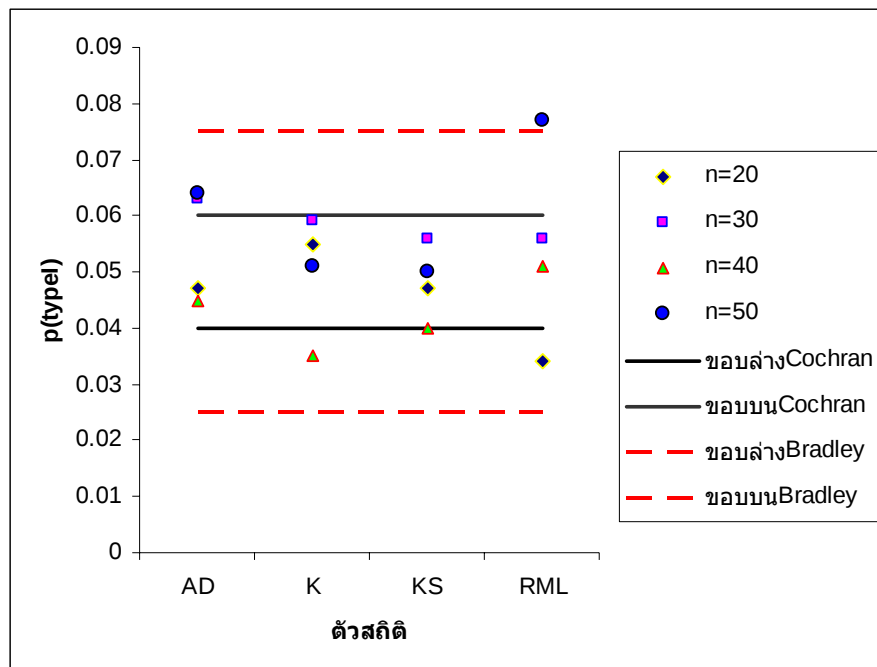
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-25 – 4-27 จำแนกตามระดับนัยสำคัญ



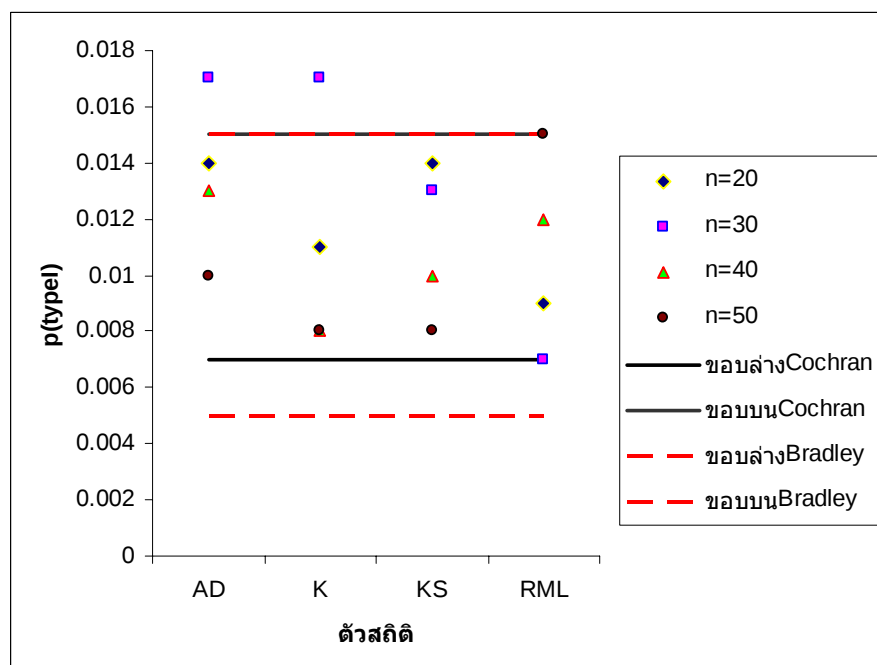
ภาพที่ 4-25 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ

AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$

ที่ $\alpha_1 = 0.1$



ภาพที่ 4-26 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-27 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.01$

ตารางที่ 4-11 แสดงค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$

n	α_1	AD	K	KS	RML
20	0.1	0.100	0.104	0.082	0.100
	0.05	0.049	0.043	0.042	0.037*
	0.01	0.009	0.003**	0.002**	0.006*
30	0.1	0.126*	0.117	0.104	0.100
	0.05	0.063*	0.057	0.059	0.052
	0.01	0.013	0.014	0.012	0.011
40	0.1	0.110	0.113	0.099	0.104
	0.05	0.063*	0.059	0.049	0.054
	0.01	0.013	0.013	0.015	0.011
50	0.1	0.112	0.104	0.094	0.091
	0.05	0.058	0.058	0.047	0.048
	0.01	0.009	0.006*	0.012	0.013

* กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran

** กรณีที่ตัวสถิติทดสอบไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ เมื่อใช้เกณฑ์ของ Cochran และเกณฑ์ของ Bradley

จากตารางที่ 4.-11 ผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 จากผลการทดลองของตัวสถิติทดสอบทั้ง 4 ตัว กับความน่าจะเป็นของการเกิดความคลาดเคลื่อนประเภทที่ 1 ที่กำหนดไว้ โดยใช้เกณฑ์ดังต่อไปนี้

1) เกณฑ์ของ Cochran

- ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 30$ ระดับนัยสำคัญ 0.1 และ 0.05 และในกรณีที่ $n = 40$ ระดับนัยสำคัญ 0.05

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01 และในกรณีที่ $n = 50$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ RML ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01 และ 0.05

2) เกณฑ์ของ Bradley

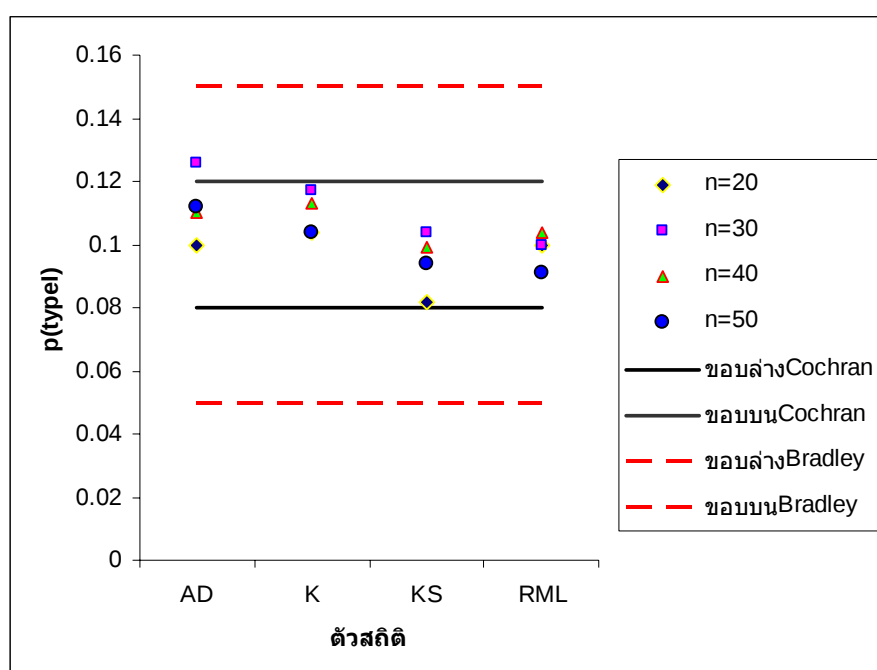
- ตัวสถิติทดสอบ AD สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

- ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

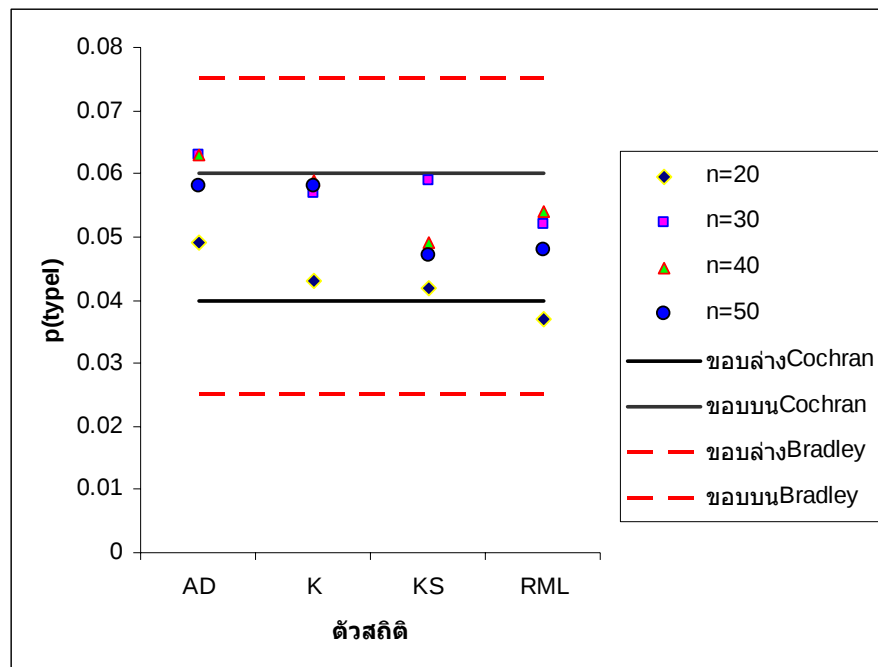
- ตัวสถิติทดสอบ KS ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ ในกรณีที่ $n = 20$ ระดับนัยสำคัญ 0.01

- ตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ครบทุกกรณี

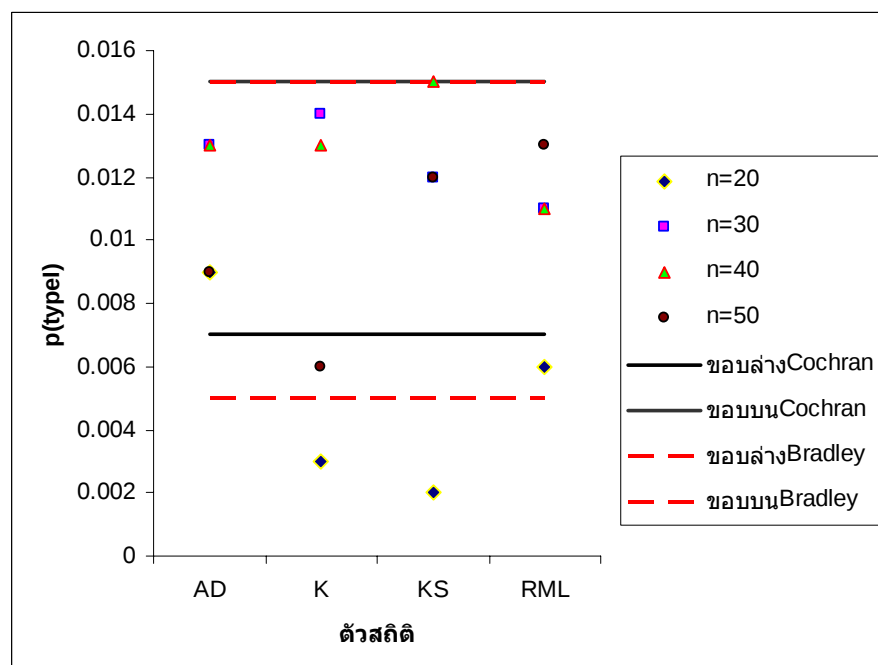
ซึ่งจะแสดงผลการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 โดยเกณฑ์ทั้ง 2 เกณฑ์ ในภาพที่ 4-28 - 4-30 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4.28 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4.29 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4.30 แสดงการเปรียบเทียบค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของสถิติ AD, K, KS, RML เมื่อ H_0 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.01$

4.2 การเปรียบเทียบอำนาจการทดสอบ

ผลการวิจัยสำหรับตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงลอกนอร์มัล

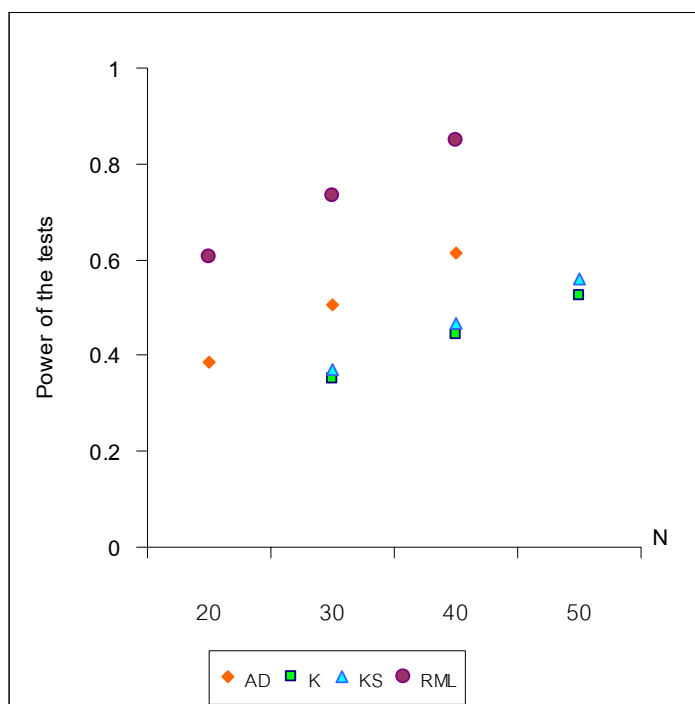
การนำเสนอค่าอำนาจการทดสอบจากการทดลองของสถิติทดสอบทั้ง 4 ตัว สำหรับตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงลอกนอร์มัล ได้แสดงเป็นตารางที่ 4-12 – 4-16

ตารางที่ 4-12 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ

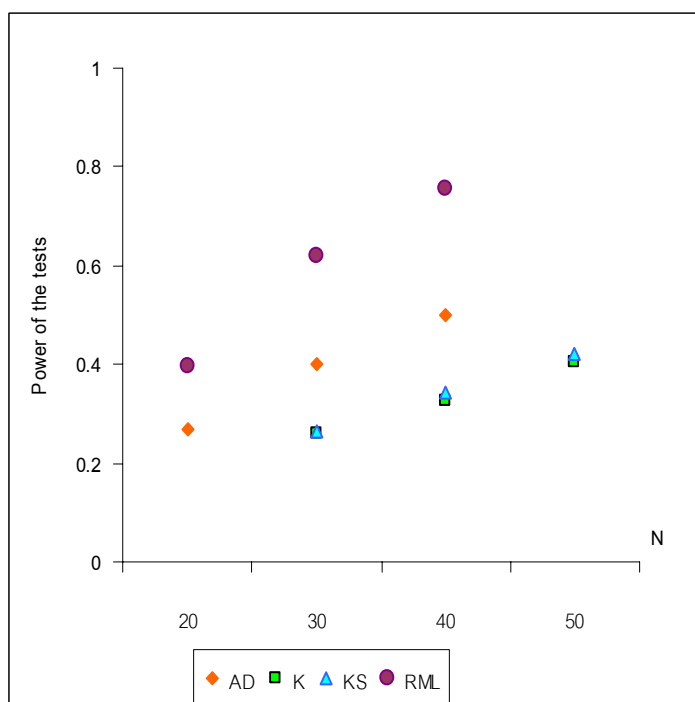
เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 0.5$

α_1	n	20	30	40	50
0.1	AD	0.385	0.504	0.613	-
	K	-	0.351	0.444	0.527
	KS	-	0.372	0.468	0.56
	RML	0.605	0.732	0.85	-
0.05	AD	0.27	0.4	0.499	-
	K	-	0.262	0.328	0.407
	KS	-	0.263	0.345	0.421
	RML	0.398	0.618	0.757	-
0.01	AD	0.141	-	0.303	-
	K	-	0.112	0.162	0.188
	KS	-	0.113	0.165	0.197
	RML	0.201	0.36	0.502	-

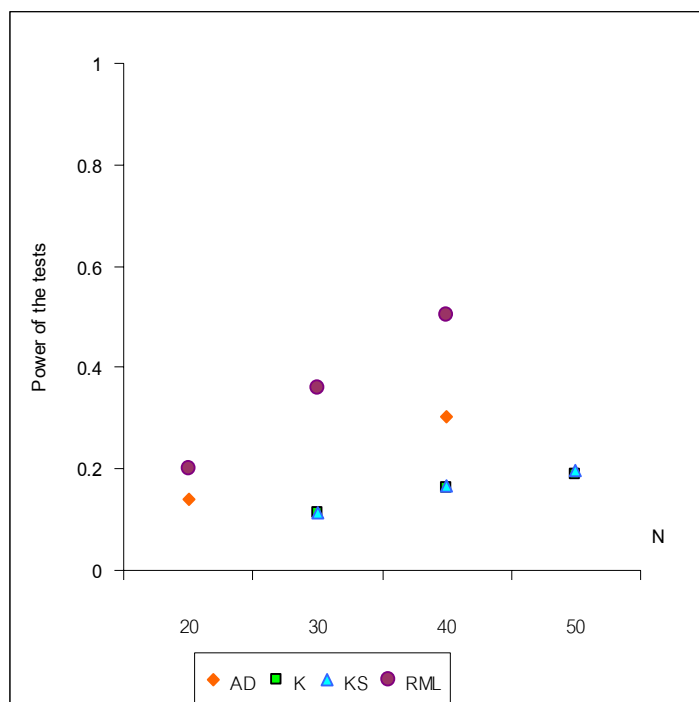
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของของสถิติ AD, K, KS, RML ในภาพที่ 4-31- 4-33 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-31 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลส์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1=0.1$



ภาพที่ 4-32 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลส์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1=0.05$



ภาพที่ 4-33 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=0.5$ ที่ $\alpha_1=0.01$

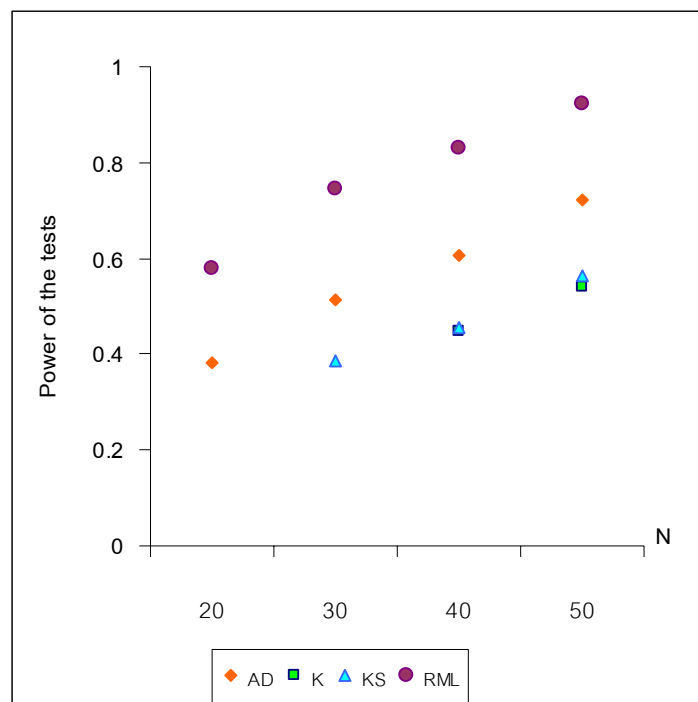
จากตารางที่ 4-12 และภาพที่ 4-30 – 4-33 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในทุกระดับนัยสำคัญ กรณีที่ $n = 20$, $n = 30$, $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD แต่ในกรณีที่ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ KS มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ K
- 2) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 3) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ K มีค่าต่ำสุดทุก ๆ กรณี
- 4) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 5) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบมีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

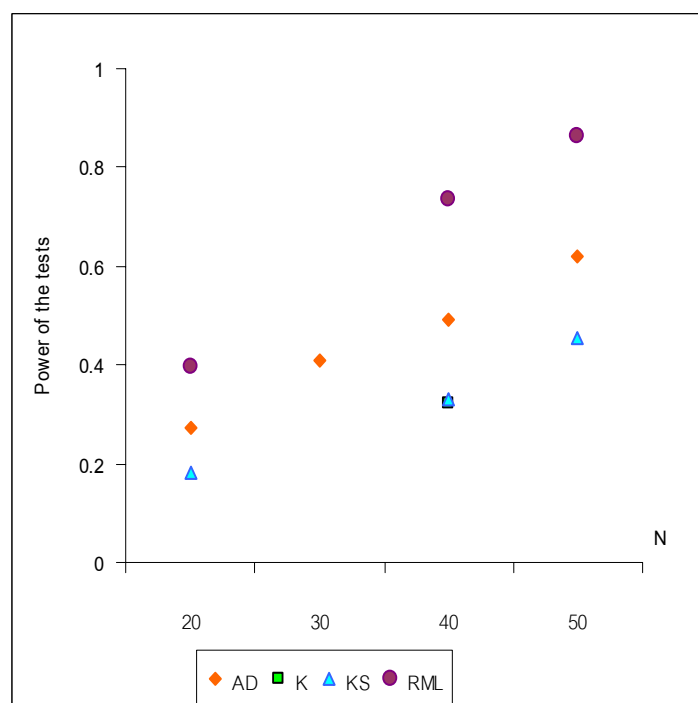
ตารางที่ 4-13 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 1$

α_1	n	20	30	40	50
0.1	AD	0.384	0.514	0.607	0.721
	K	-	-	0.446	0.539
	KS	-	0.385	0.456	0.564
	RML	0.579	0.746	0.83	0.923
0.05	AD	0.274	0.408	0.492	0.618
	K	-	-	0.324	-
	KS	0.183	-	0.329	0.454
	RML	0.398	-	0.735	0.864
0.01	AD	-	0.217	0.295	0.396
	K	0.08	-	0.149	0.224
	KS	0.068	-	0.148	0.235
	RML	0.21	0.383	0.517	0.653

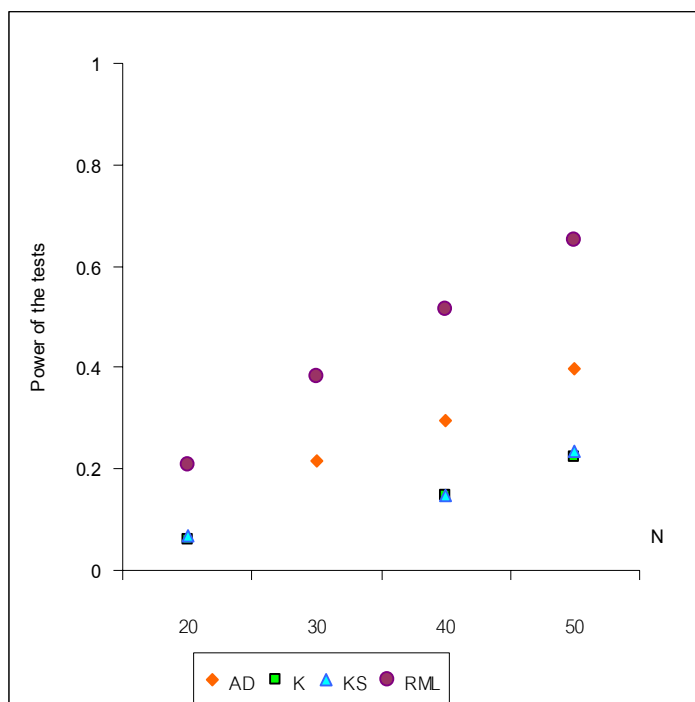
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-34 - 4-36 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-34 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1$ ที่ $\alpha_1=0.1$



ภาพที่ 4-35 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1$ ที่ $\alpha_1=0.05$



ภาพที่ 4-36 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1$ ที่ $\alpha_1 = 0.01$

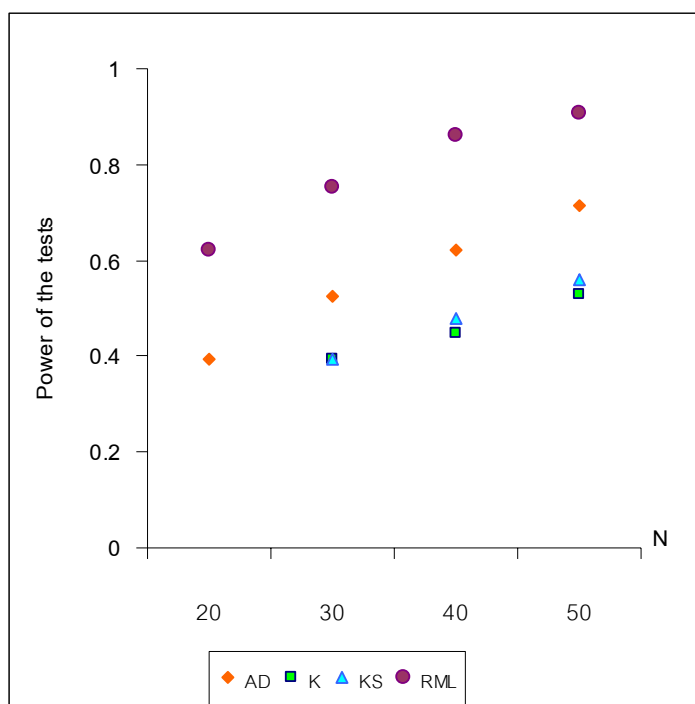
จากตารางที่ 4-13 และภาพที่ 4-34 – 4-36 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ที่ระดับนัยสำคัญ 0.1 และ 0.01 ในทุก ๆ ค่า n พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 2) ที่ระดับนัยสำคัญ 0.05 ในกรณีที่ $n = 30$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด แต่ในกรณี n อื่น ๆ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 3) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 4) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ K มีค่าต่ำสุดทุก ๆ กรณี
- 5) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 6) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบมีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

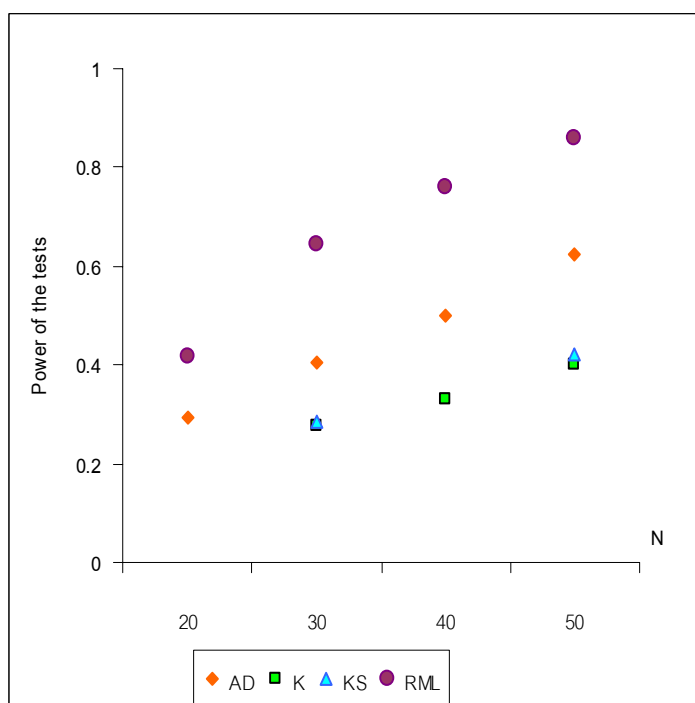
ตารางที่ 4-14 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1.5$

α_1	n	20	30	40	50
0.1	AD	0.395	0.527	0.621	0.715
	K	-	0.392	0.448	0.529
	KS	-	0.395	0.477	0.561
	RML	0.622	0.751	0.862	0.909
0.05	AD	0.292	0.405	0.5	0.626
	K	-	0.275	0.33	0.402
	KS	-	0.284	-	0.422
	RML	0.416	0.644	0.76	0.858
0.01	AD	0.123	0.227	0.297	0.402
	K	0.069	0.117	-	0.224
	KS	0.073	-	-	0.233
	RML	0.208	0.395	0.525	0.661

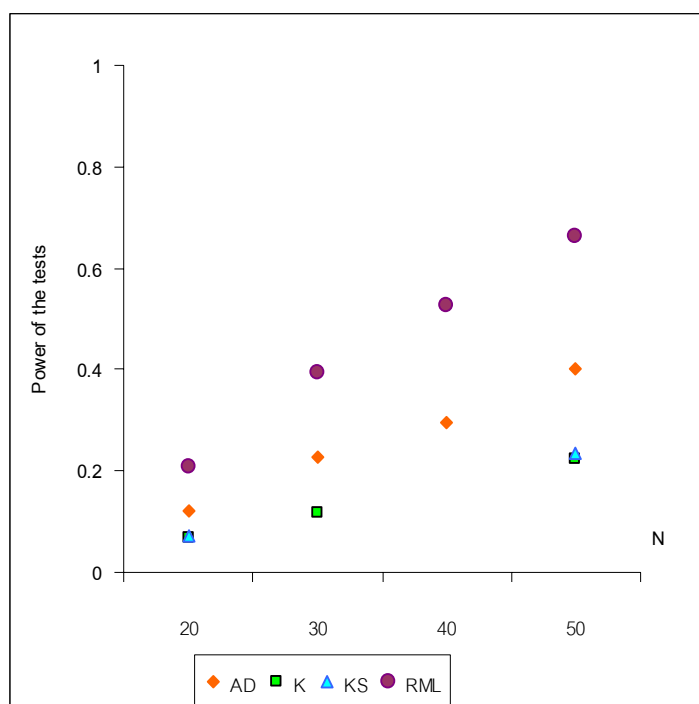
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4.37- 4.39 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-37 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1.5$ ที่ $\alpha_1=0.1$



ภาพที่ 4-38 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1.5$ ที่ $\alpha_1=0.05$



ภาพที่ 4-39 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=1.5$ ที่ $\alpha_1 = 0.01$

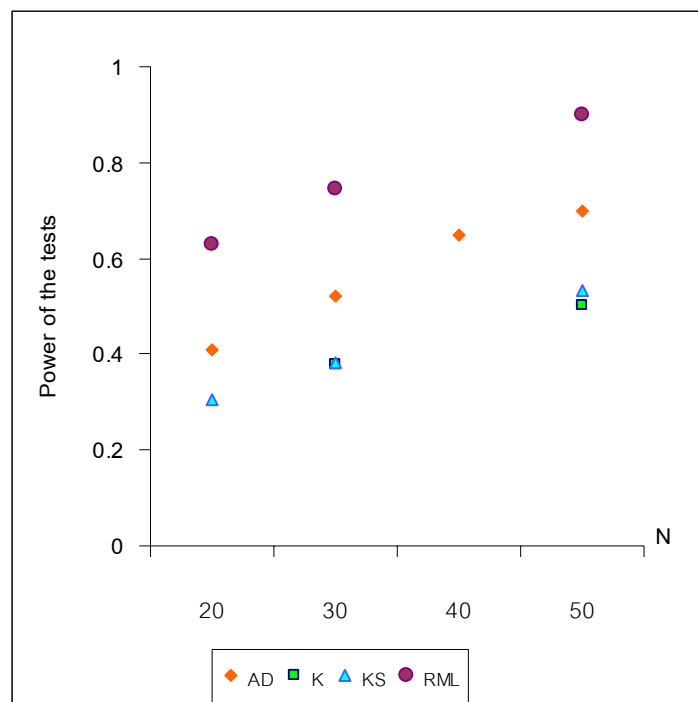
จากตารางที่ 4-14 และภาพที่ 4-37 – 4-39 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) กรณีที่วิเคราะห์ข้อมูลทั้งหมด พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD
- 2) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 3) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ K มีค่าต่ำสุดทุก ๆ กรณี
- 4) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 5) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบมีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

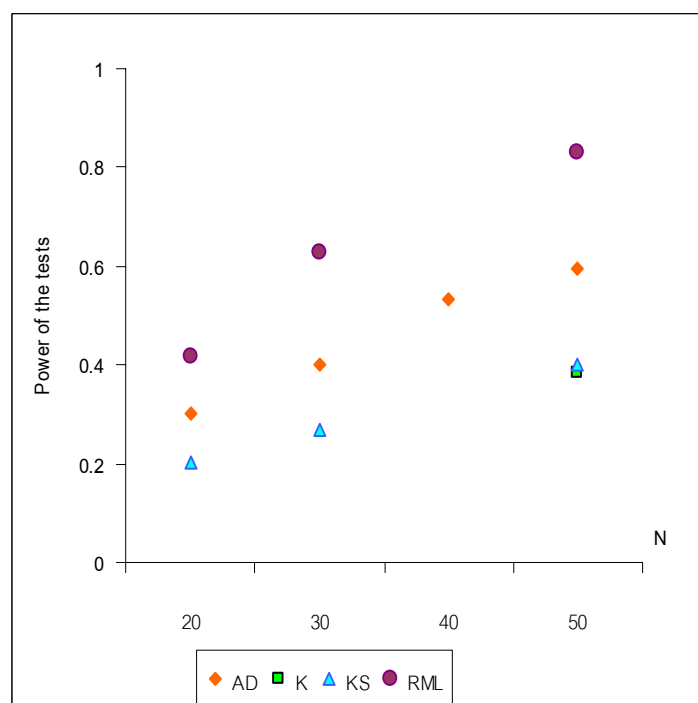
ตารางที่ 4-15 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 2$

α_1	n	20	30	40	50
0.1	AD	0.411	0.521	0.649	0.698
	K	-	0.378	-	0.501
	KS	0.304	0.383	-	0.534
	RML	0.629	0.747	-	0.901
0.05	AD	0.303	0.399	0.532	0.596
	K	-	-	-	0.386
	KS	0.204	0.268	-	0.402
	RML	0.416	0.63	-	0.832
0.01	AD	0.145	0.226	0.314	0.367
	K	0.08	-	0.169	0.2
	KS	0.089	-	-	0.2
	RML	0.226	0.394	-	-

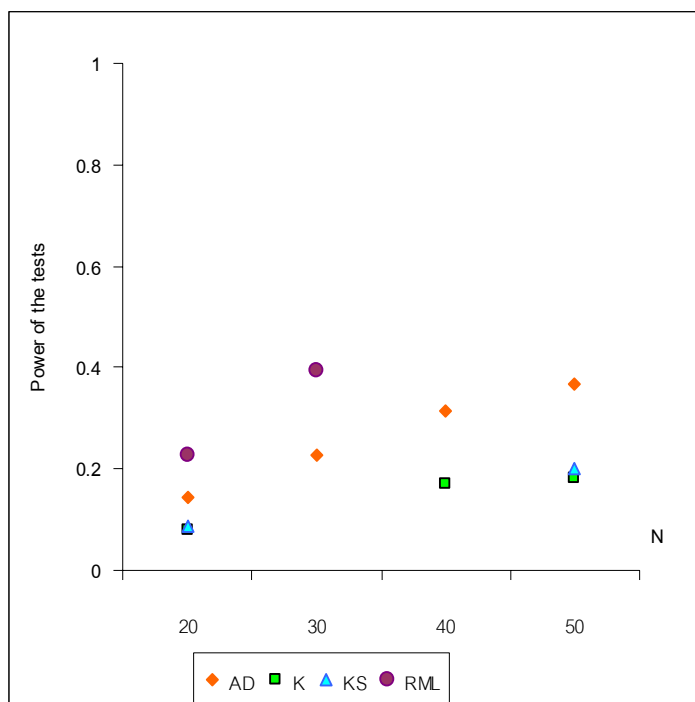
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-40 - 4-42 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-40 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=2$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-41 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=2$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-42 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=2$ ที่ $\alpha_1=0.01$

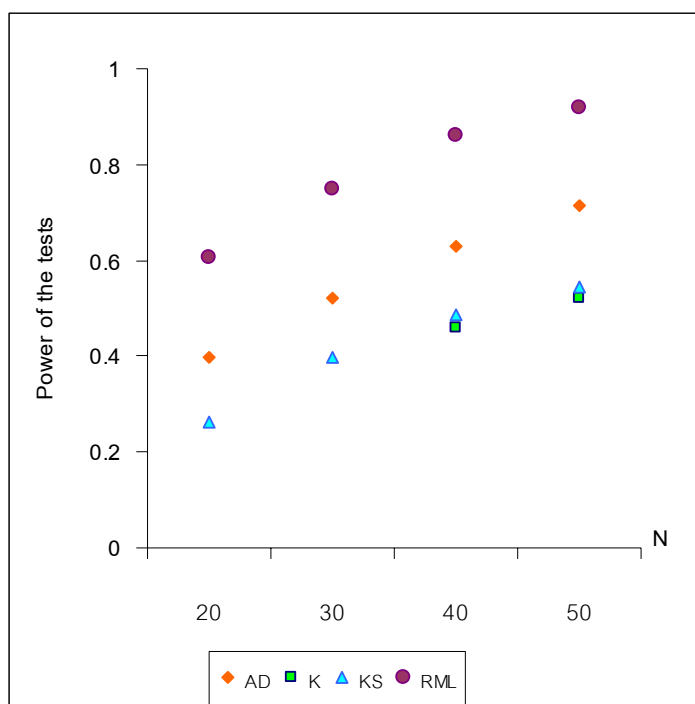
จากตารางที่ 4-15 และภาพที่ 4-40 – 4-42 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในระดับนัยสำคัญที่ 0.1 และ 0.05 กรณีที่ $n = 20$, $n = 30$, $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD แต่ในกรณีที่ $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 2) ในระดับนัยสำคัญที่ 0.01 กรณีที่ $n = 20$, $n = 30$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD แต่ในกรณีที่ $n = 40$ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 3) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ K มีค่าต่ำสุดทุก ๆ กรณี
- 4) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 5) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 6) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบมีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

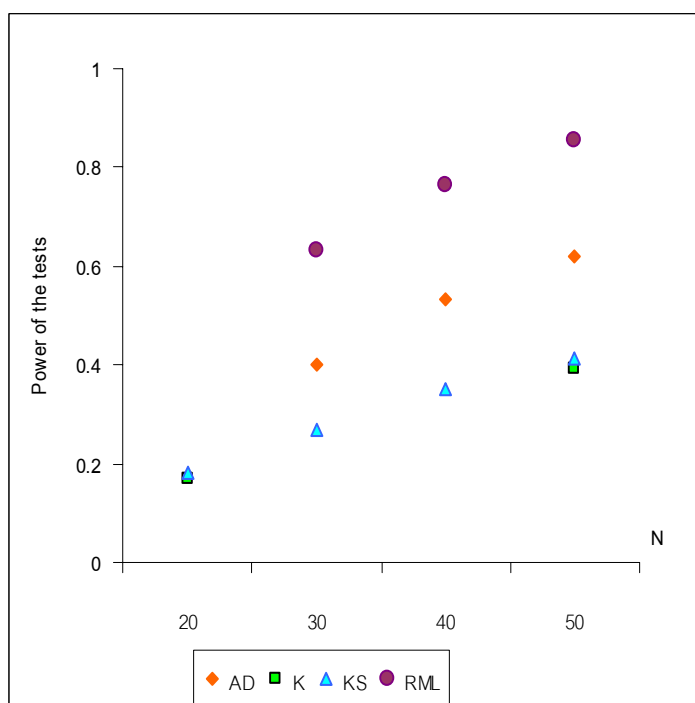
ตารางที่ 4-16 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta = 1$ และ $\alpha = 3$

α_1	n	20	30	40	50
0.1	AD	0.397	0.52	0.628	0.714
	K	-	-	0.461	0.523
	KS	0.261	0.398	0.487	0.543
	RML	0.607	0.749	0.86	0.92
0.05	AD	-	0.4	0.531	0.621
	K	0.17	-	-	0.392
	KS	0.18	0.268	0.353	0.414
	RML	-	0.633	0.764	0.856
0.01	AD	-	0.209	0.322	0.396
	K	0.062	0.11	-	0.206
	KS	0.073	0.113	-	0.211
	RML	0.208	0.383	0.553	0.625

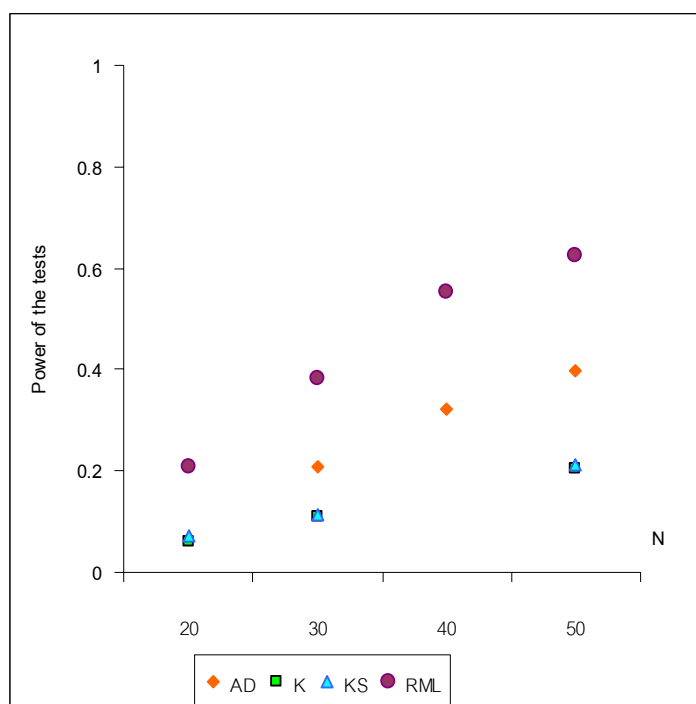
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-43 - 4-45 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-43 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=3$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-44 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=3$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-45 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงไวบูลล์ $\beta=1$ และ $\alpha=3$ ที่ $\alpha_1 = 0.01$

จากตารางที่ 4-16 และภาพที่ 4-43 - 4-45 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในระดับนัยสำคัญที่ 0.1 และ 0.01 ในทุกกรณี พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 2) ในระดับนัยสำคัญที่ 0.05 กรณีที่ $n = 30$, $n = 40$, $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 20$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ KS มีค่าสูงที่สุด
- 3) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ K มีค่าต่ำสุดทุก ๆ กรณี
- 4) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 5) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 6) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้น เมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบมีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

ผลการวิจัยสำหรับตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงไวบูลล์

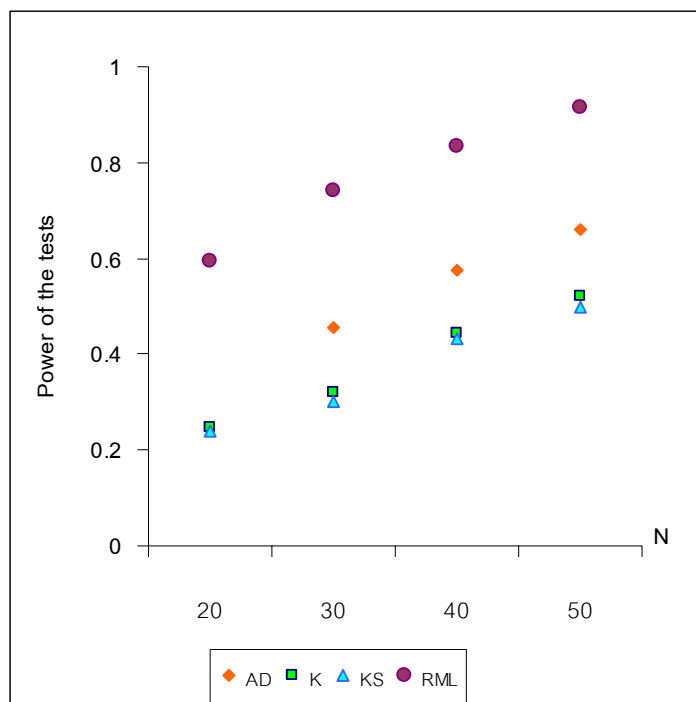
การนำเสนอค่าอำนาจการทดสอบจากการทดลองของสถิติทดสอบทั้ง 4 ตัว สำหรับตัวสถิติทดสอบที่ใช้ทดสอบการแจกแจงไวบูลล์ได้แสดงเป็นตารางที่ 4-17 – 4-21

ตารางที่ 4-17 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ

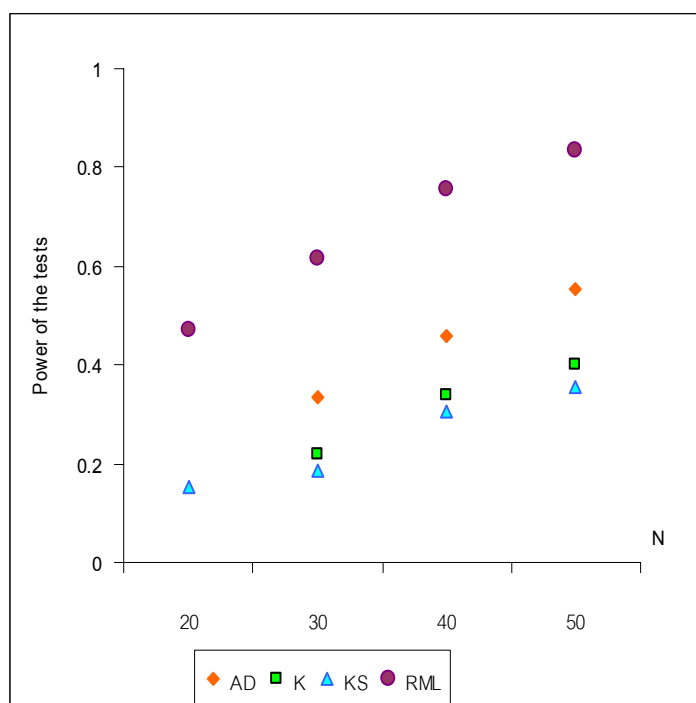
เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 0.5$

α_1	n	20	30	40	50
0.1	AD	-	0.456	0.574	0.66
	K	0.247	0.321	0.443	0.521
	KS	0.24	0.301	0.431	0.5
	RML	0.595	0.743	0.833	0.916
0.05	AD	-	0.333	0.46	0.555
	K	-	0.218	0.337	0.402
	KS	0.152	0.185	0.306	0.356
	RML	0.472	0.614	0.757	0.836
0.01	AD	0.098	0.166	0.243	-
	K	0.061	0.088	0.162	0.227
	KS	0.049	0.075	-	0.169
	RML	0.235	0.34	0.502	0.653

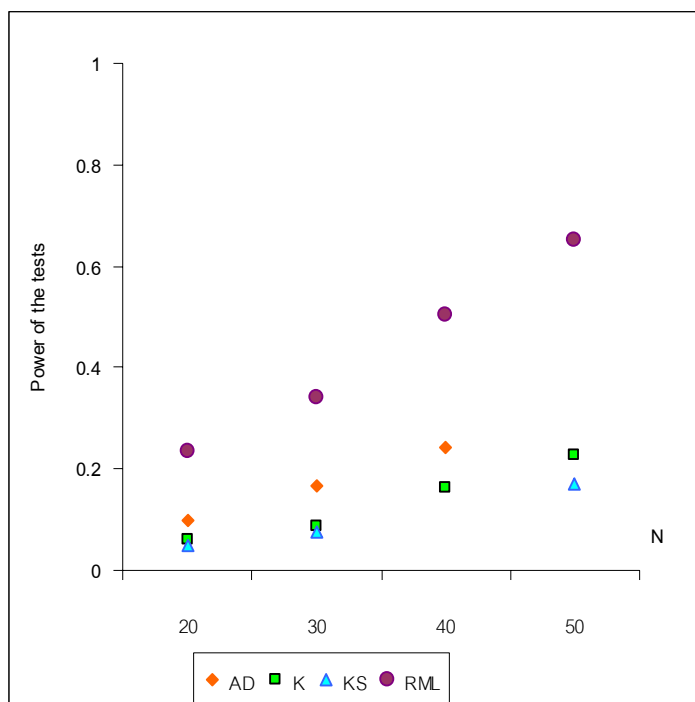
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ในภาพที่ 4-46 – 4-48 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-46 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-47 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอกนอร์มัล $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.01$



ภาพที่ 4-48 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการมัล $\mu = 1$ และ $\sigma = 0.5$ ที่ $\alpha_1 = 0.01$

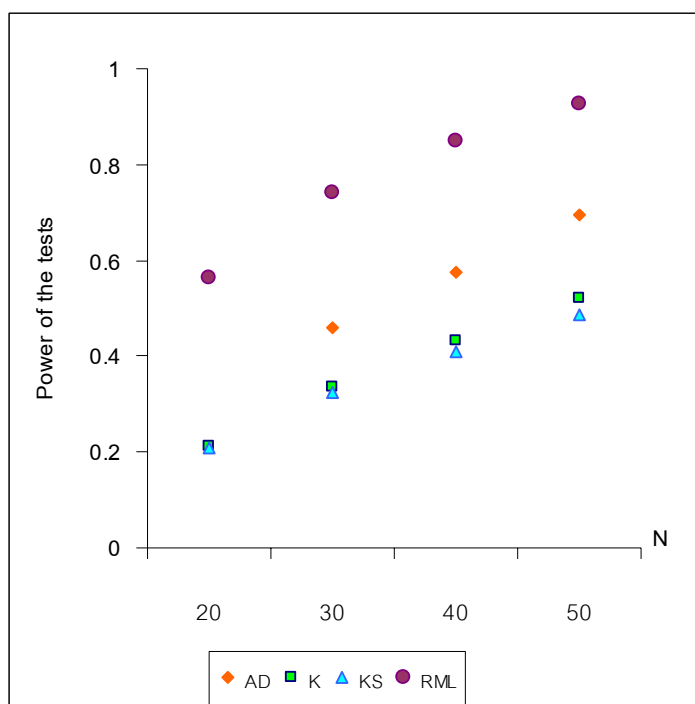
จากตารางที่ 4-17 และภาพที่ 4-46 – 4-48 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในทุกระดับนัยสำคัญ และในทุกกรณี พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 2) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 3) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ KS มีค่าต่ำสุดทุก ๆ กรณี
- 4) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 5) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบมีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

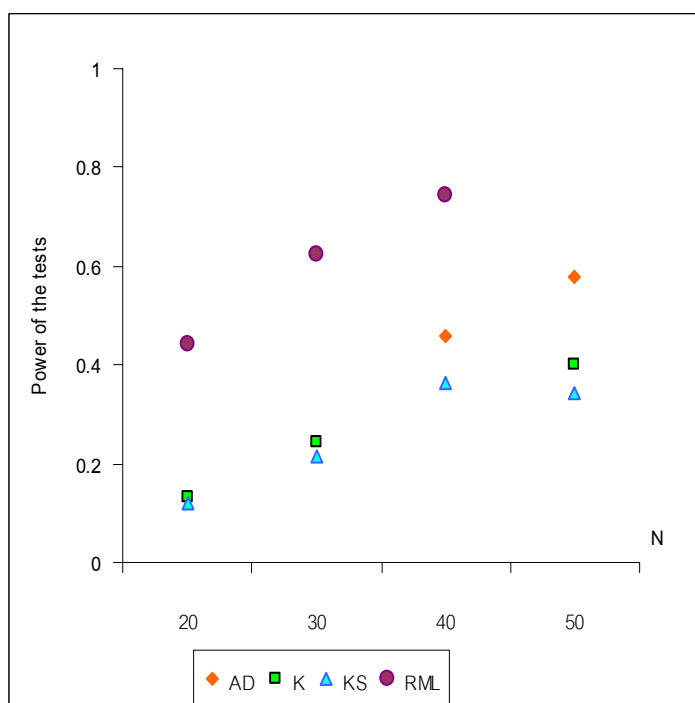
ตารางที่ 4-18 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$

α_1	n	20	30	40	50
0.1	AD	-	0.458	0.576	0.695
	K	0.214	0.335	0.431	0.521
	KS	0.207	0.326	0.41	0.488
	RML	0.565	0.741	0.848	0.925
0.05	AD	-	-	0.46	0.577
	K	0.132	0.244	-	0.401
	KS	0.119	0.213	0.365	0.344
	RML	0.442	0.625	0.743	-
0.01	AD	0.085	0.165	-	0.349
	K	-	0.092	-	-
	KS	0.027	0.078	0.115	0.163
	RML	-	0.347	0.496	-

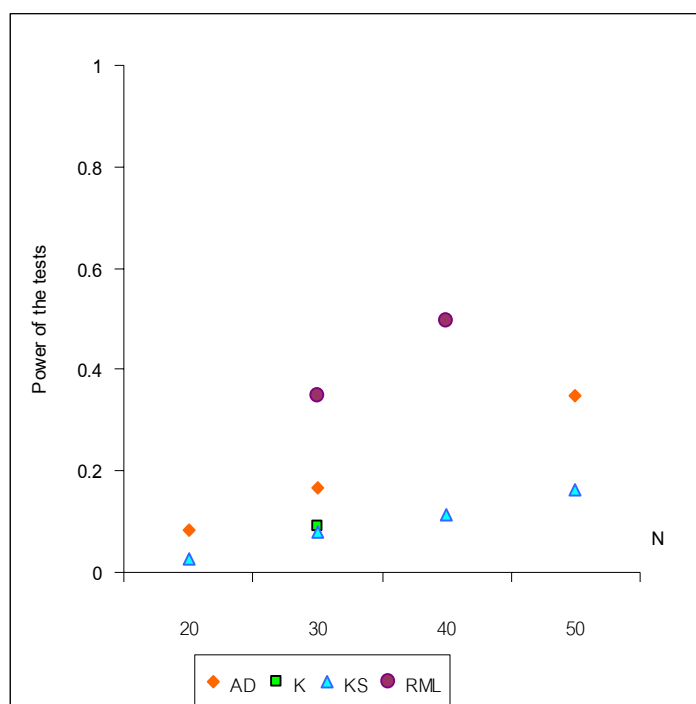
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-49 – 4-51 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-49 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-50 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-51 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงลอการมัล $\mu=1$ และ $\sigma=1$ ที่ $\alpha_1 = 0.01$

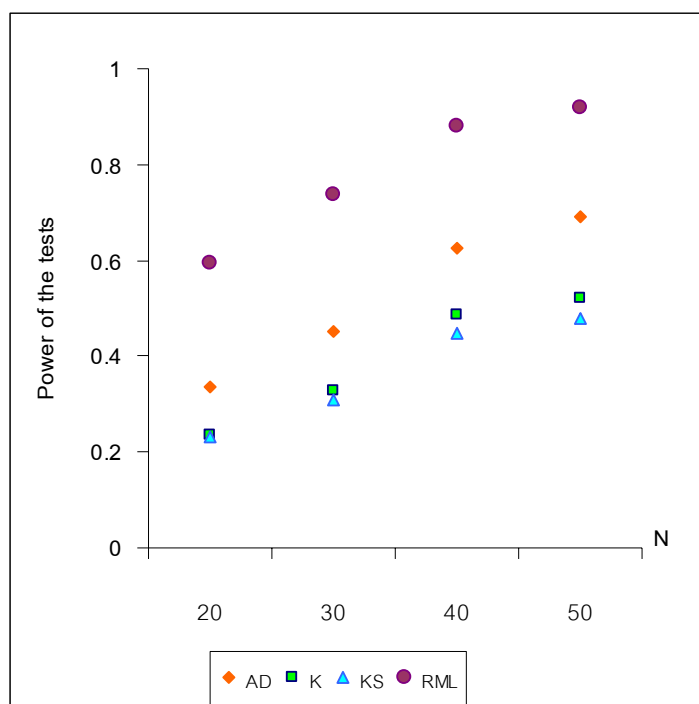
จากตารางที่ 4-18 และภาพที่ 4-49 – 4-51 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในระดับนัยสำคัญที่ 0.1 ในทุกกรณี พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 2) ในระดับนัยสำคัญที่ 0.05 กรณีที่ $n = 20$, $n = 30$, $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 3) ในระดับนัยสำคัญที่ 0.01 กรณีที่ $n = 30$, $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 20$ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 4) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ KS มีค่าต่ำสุดทุก ๆ กรณี
- 5) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 6) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 7) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบ มีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

ตารางที่ 4-19 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu=1$ และ $\sigma=1.5$

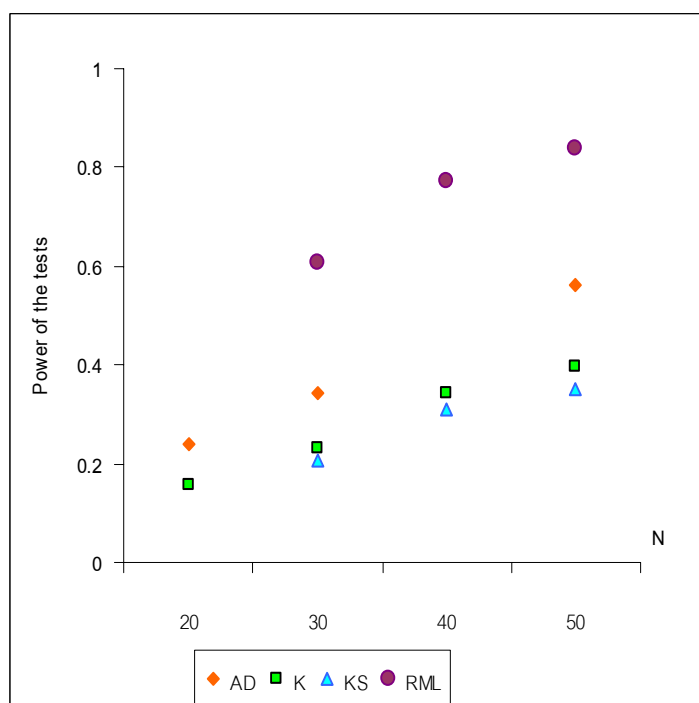
α_1	n	20	30	40	50
0.1	AD	0.334	0.452	0.625	0.69
	K	0.234	0.329	0.487	0.521
	KS	0.232	0.309	0.446	0.479
	RML	0.593	0.737	0.879	0.92
0.05	AD	0.241	0.343	-	0.562
	K	0.157	0.233	0.341	0.395
	KS	-	0.205	0.308	0.35
	RML	-	0.609	0.774	0.839
0.01	AD	-	0.167	0.271	0.347
	K	0.068	0.098	0.17	0.217
	KS	-	0.083	0.122	0.17
	RML	-	0.344	0.553	-

ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-52 – 4-54 จำแนกตามระดับนัยสำคัญ



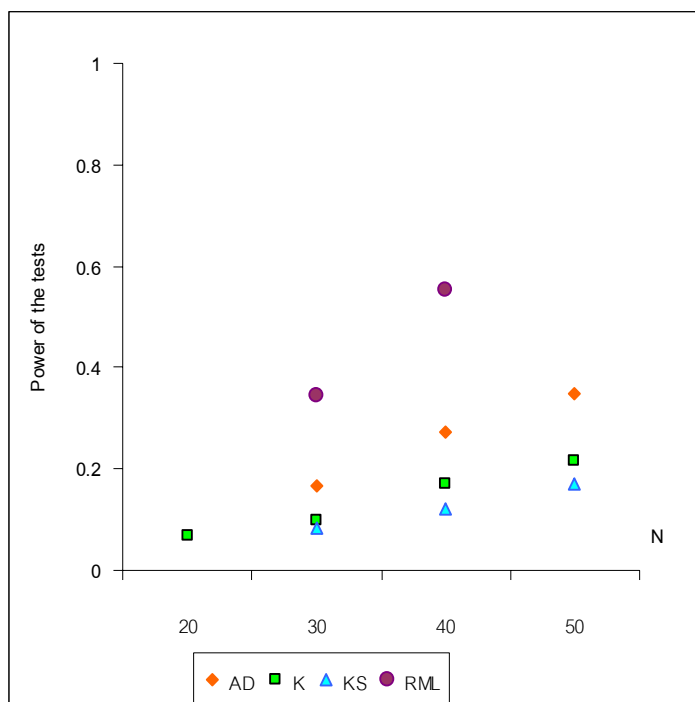
ภาพที่ 4-52 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงลอการมัล $\mu=1$ และ $\sigma=1.5$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-53 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงลอการมัล $\mu=1$ และ $\sigma=1.5$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-54 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 1.5$ ที่ $\alpha_1 = 0.01$

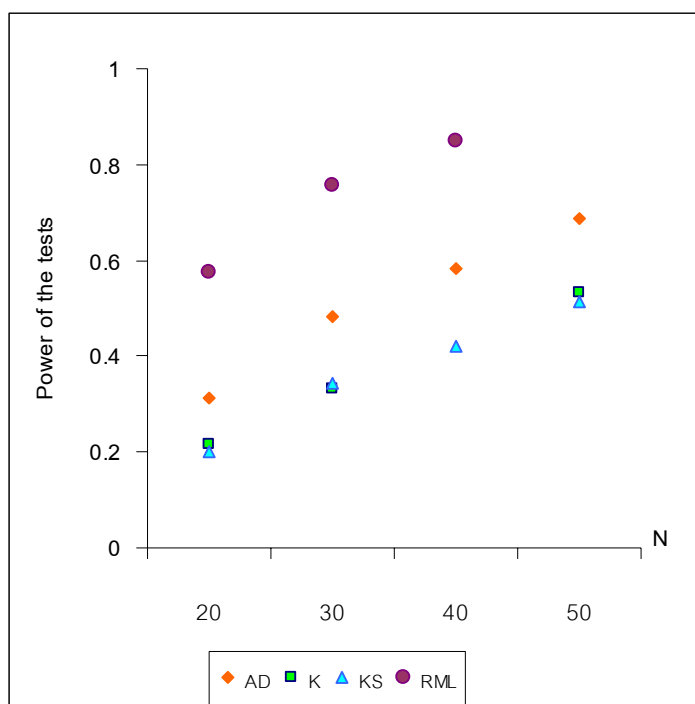
จากตารางที่ 4-19 และภาพที่ 4-52 – 4-54 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในระดับนัยสำคัญที่ 0.1 ในทุกกรณี พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 2) ในระดับนัยสำคัญที่ 0.05 กรณีที่ $n = 30$, $n = 40$, $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 20$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 3) ในระดับนัยสำคัญที่ 0.01 กรณีที่ $n = 30$, $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด ในกรณีที่ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด และในกรณีที่ $n = 20$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ K มีค่าสูงที่สุด
- 4) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ KS มีค่าต่ำสุดทุก ๆ กรณี
- 5) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 6) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้นทุกระดับนัยสำคัญ
- 7) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบ มีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

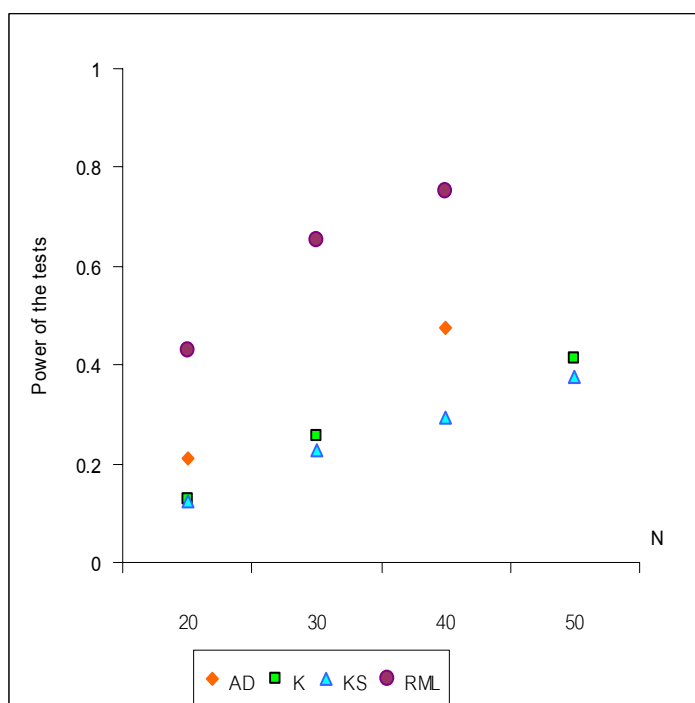
ตารางที่ 4-20 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$

α_1	n	20	30	40	50
0.1	AD	0.314	0.483	0.583	0.686
	K	0.218	0.331	-	0.532
	KS	0.2	0.343	0.419	0.514
	RML	0.574	0.758	0.848	-
0.05	AD	0.209	-	0.474	-
	K	0.129	0.257	-	0.415
	KS	0.122	0.228	0.292	0.378
	RML	0.429	0.653	0.754	-
0.01	AD	0.086	-	0.269	0.359
	K	0.052	-	0.155	0.237
	KS	0.043	0.089	0.144	0.193
	RML	0.201	0.377	0.522	0.667

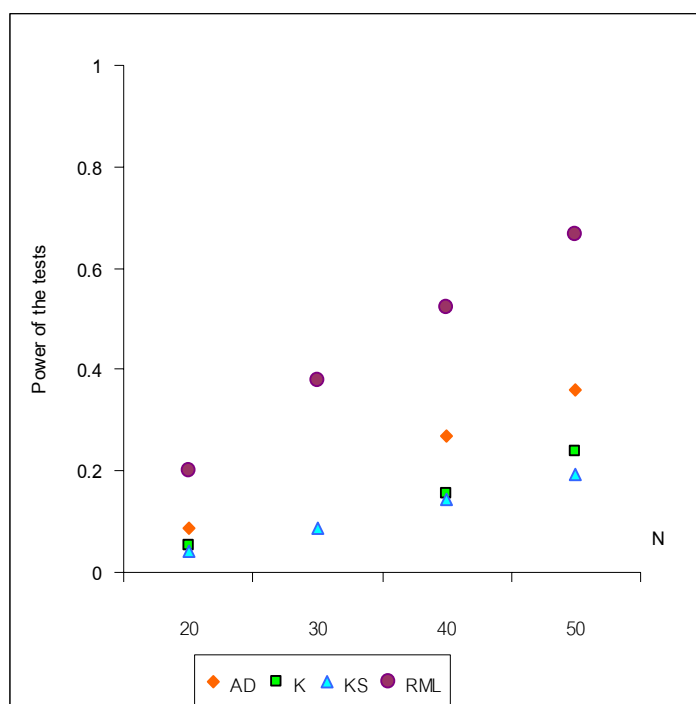
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-55 – 4-57 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-55 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-56 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-57 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ

H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 2$ ที่ $\alpha_1 = 0.01$

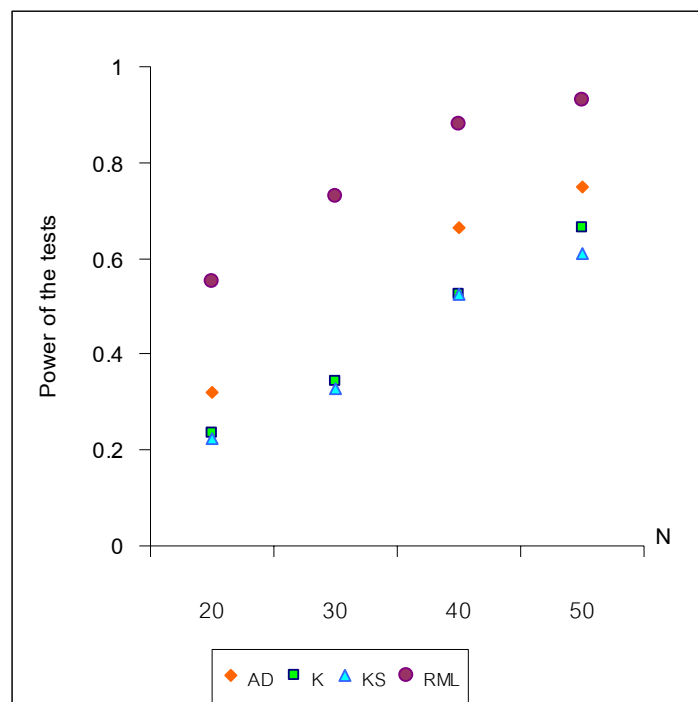
จากตารางที่ 4-20 และภาพที่ 4-55 – 4-57 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในระดับนัยสำคัญที่ 0.1 กรณีที่ $n = 20$, $n = 30$, $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 2) ในระดับนัยสำคัญที่ 0.1 กรณีที่ $n = 20$, $n = 30$, $n = 40$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ K มีค่าสูงที่สุด
- 3) ในระดับนัยสำคัญที่ 0.01 ทุกกรณี พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 4) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ KS มีค่าต่ำสุดทุก ๆ กรณี
- 5) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 6) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 7) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบ มีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

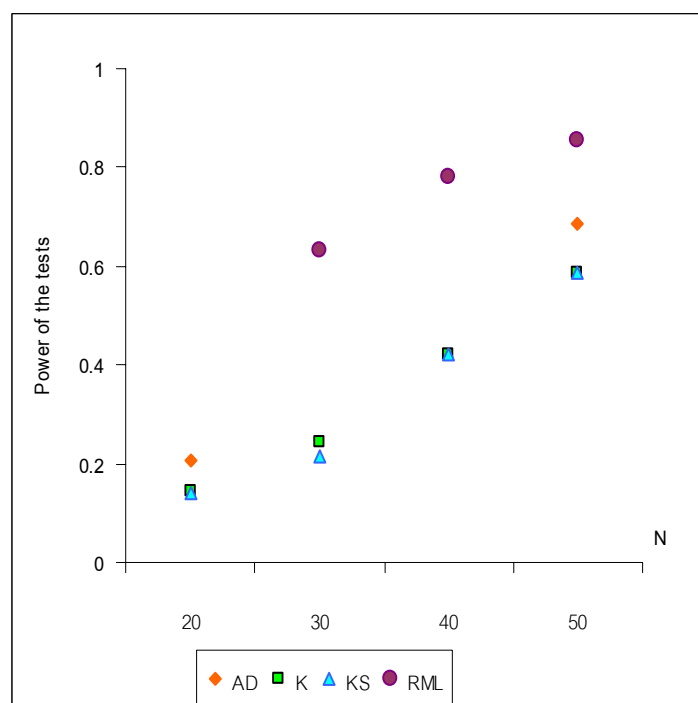
ตารางที่ 4-21 แสดงค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML กระทำซ้ำ 1000 รอบ
เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$

α_1	n	20	30	40	50
0.1	AD	0.321	-	0.666	0.749
	K	0.236	0.345	0.525	0.665
	KS	0.223	0.329	0.525	0.61
	RML	0.554	0.73	0.881	0.932
0.05	AD	0.205	-	-	0.686
	K	0.143	0.244	0.423	0.588
	KS	0.141	0.213	0.423	0.585
	RML	-	0.631	0.779	0.856
0.01	AD	0.088	0.177	0.469	0.562
	K	-	0.106	0.265	-
	KS	-	0.084	0.265	0.478
	RML	-	0.358	0.604	0.725

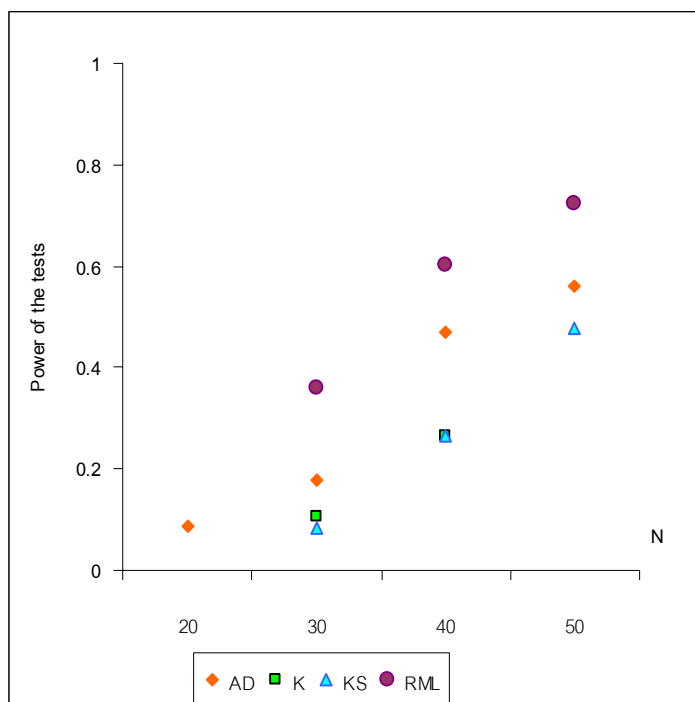
ซึ่งจะแสดงผลการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML ใน
ภาพที่ 4-58 – 4-60 จำแนกตามระดับนัยสำคัญ



ภาพที่ 4-58 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.1$



ภาพที่ 4-59 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการิธึม $\mu = 1$ และ $\sigma = 3$ ที่ $\alpha_1 = 0.05$



ภาพที่ 4-60 แสดงการเปรียบเทียบค่าอำนาจการทดสอบของสถิติ AD, K, KS, RML เมื่อ H_1 : ประชากรมีการแจกแจงลอการมัล $\mu=1$ และ $\sigma=3$ ที่ $\alpha_1 = 0.01$

จากตารางที่ 4-21 และภาพที่ 4-58 – 4-60 เราสามารถสรุปผลการวิเคราะห์ได้ดังนี้

- 1) ในระดับนัยสำคัญที่ 0.1 ทุกกรณี พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด
- 2) ในระดับนัยสำคัญที่ 0.05 และ 0.01 กรณีที่ $n = 30, n = 40, n = 50$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด แต่ในกรณีที่ $n = 20$ พบว่าค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD มีค่าสูงที่สุด
- 3) ในทุกระดับนัยสำคัญตัวสถิติทดสอบ KS มีค่าต่ำสุดทุก ๆ กรณี
- 4) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 5) เมื่อ n เพิ่มขึ้นค่าอำนาจการทดสอบของตัวสถิติทดสอบทุกตัวจะเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 6) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้นเมื่อเพิ่มค่าระดับนัยสำคัญในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบ มีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

สรุปผลการเปรียบเทียบอำนาจการทดสอบของตัวสถิติที่ใช้ทดสอบการแจกแจงแบบลอกนอร์มัล และตัวสถิติที่ใช้ทดสอบการแจกแจงแบบไวบูลล์

ผลการทดสอบโดยใช้ตัวสถิติทั้ง 4 ตัว คือ AD, K, KS และ RML ทั้งที่ใช้ทดสอบการแจกแจงแบบลอกนอร์มัล และการแจกแจงแบบไวบูลล์ เมื่อขนาดตัวอย่างเท่ากับ 20, 30, 40 และ 50 ณ ระดับนัยสำคัญ 0.1, 0.05 และ 0.01 เป็นดังนี้

- 1) กรณีที่วิเคราะห์ข้อมูลทั้งหมด พบว่าทุก ๆ กรณีที่สามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้ค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD ทุกระดับนัยสำคัญ
- 2) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน
- 3) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K มีค่าต่ำสุด สำหรับตัวสถิติทดสอบที่ใช้การแจกแจงแบบลอกนอร์มัล
- 4) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ KS มีค่าต่ำสุด สำหรับตัวสถิติทดสอบที่ใช้การแจกแจงแบบไวบูลล์
- 5) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD, K, KS และ RML จะสูงขึ้น เมื่อขนาดตัวอย่างเพิ่มขึ้น ทุกระดับนัยสำคัญ
- 6) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้น เมื่อระดับนัยสำคัญเพิ่มขึ้น ในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบ มีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

บทที่ 5

สรุปผลการวิจัย

ในการวิจัยครั้งนี้เป็นการวิจัยเชิงทดลองเพื่อทำการทดสอบภาวะสภาวะรูปดี (Goodness of Fit Test : GOF) สำหรับการจำแนกการแจกแบบลอการิทึม (Log-normal Distribution) และการแจกแบบไวบูลล์ (Weibull Distribution) ซึ่งสถิติทดสอบที่ใช้ในการวิจัย ได้แก่ Anderson Darling (AD) Test Statistic, Kolmogorov –Smirnov (KS) Test Statistic, Kuiper (K) Test Statistic และ Ratio of Maximum Likelihoods (RML) Test Statistic โดยสถานการณ์ต่างๆ ที่กำหนดขึ้นดังนี้

1. กำหนดลักษณะของข้อมูล มีการแจกแจง 2 การแจกแจง ดังนี้

1.1 การแจกแจงแบบลอการิทึม (Log-normal Distribution) ที่พารามิเตอร์ ดังนี้

$$\mu = 1 \text{ และ } \sigma = 0.5$$

$$\mu = 1 \text{ และ } \sigma = 1$$

$$\mu = 1 \text{ และ } \sigma = 1.5$$

$$\mu = 1 \text{ และ } \sigma = 2$$

$$\mu = 1 \text{ และ } \sigma = 3$$

1.2 การแจกแจงแบบไวบูลล์ (Weibull Distribution) ที่พารามิเตอร์ ดังนี้

$$\beta = 1 \text{ และ } \alpha = 0.5$$

$$\beta = 1 \text{ และ } \alpha = 1$$

$$\beta = 1 \text{ และ } \alpha = 1.5$$

$$\beta = 1 \text{ และ } \alpha = 2$$

$$\beta = 1 \text{ และ } \alpha = 3$$

2. กำหนดระดับนัยสำคัญ (α_1) 3 ระดับ ได้แก่ 0.01, 0.05, 0.10

3. ขนาดตัวอย่างที่ใช้ในการศึกษา ได้แก่ 20, 30, 40, 50

การวิจัยครั้งนี้ได้จำลองข้อมูลให้มีสถานการณ์ตามที่กำหนดข้างต้น โดยใช้โปรแกรม R และการจำลองข้อมูลซ้ำๆ กัน 1,000 ครั้ง ในแต่ละสถานการณ์

5.1 สรุปผลการวิจัย

การสรุปผลว่า ตัวสถิติทดสอบใดมีความเหมาะสม สำหรับการจำแนกการแจกแบบลอการิทึม และการแจกแบบไวบูลล์ พิจารณาความสามารถในการควบคุมความน่าจะเป็นของ

ความคลาดเคลื่อนประเภทที่ 1 จากการทดลองเป็นอันดับแรก แล้วจึงพิจารณาอำนาจการทดสอบของตัวสถิติทดสอบเป็นอันดับต่อไป ซึ่งผลสรุปทั้ง 2 ขั้นตอนเป็นดังนี้

5.1.1 ผลการเปรียบเทียบความสามารถในการควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ของตัวสถิติทดสอบ 4 ตัว ณ ระดับนัยสำคัญ 0.1, 0.05 และ 0.01 เป็นดังนี้

1) เกณฑ์ของ Cochran

ตัวสถิติทดสอบ K ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 มากที่สุด คือ 34 สถานการณ์ จากสถานการณ์ทั้งหมดในการวิจัยครั้งนี้ 120 สถานการณ์ หรือคิดเป็น 28.33% ของการทดลองทั้งหมด และตัวสถิติทดสอบ RML สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ดีที่สุด มีเพียง 19 สถานการณ์ จากสถานการณ์ทั้งหมดในการวิจัยครั้งนี้ 120 สถานการณ์ ที่ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ หรือคิดเป็น 16 % ของการทดลองทั้งหมด

2) เกณฑ์ของ Bladley

ตัวสถิติทดสอบ AD ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 มากที่สุด คือ 7 สถานการณ์ จากสถานการณ์ทั้งหมดในการวิจัยครั้งนี้ 120 สถานการณ์ หรือ คิดเป็น 5.83% ของการทดลองทั้งหมด และตัวสถิติทดสอบ KS สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ดีที่สุด มีเพียง 4 สถานการณ์ จากสถานการณ์ทั้งหมดในการวิจัยครั้งนี้ 120 สถานการณ์ ที่ไม่สามารถควบคุมความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้ หรือ คิดเป็น 3.33% ของการทดลองทั้งหมด

5.1.2 ผลการเปรียบเทียบอำนาจการทดสอบของตัวสถิติทดสอบ 4 ตัว สำหรับการจำแนกการแจกแบบลอการิธึม และการแจกแบบไวบูลล์ ณ ระดับนัยสำคัญ 0.1, 0.05 และ 0.01 เป็นดังนี้

ผลการทดสอบโดยใช้ตัวสถิติทั้ง 4 ตัว คือ AD, K, KS และ RML ทั้งที่ใช้ทดสอบการแจกแบบลอการิธึม และการแจกแบบไวบูลล์ เมื่อขนาดตัวอย่างเท่ากับ 20, 30, 40 และ 50 ณ ระดับนัยสำคัญ 0.1, 0.05 และ 0.01 เป็นดังนี้

1) กรณีที่วิเคราะห์ข้อมูลทั้งหมด พบว่าทุก ๆ กรณีที่สามารถควบคุมความคลาดเคลื่อนประเภทที่ 1 ได้ค่าอำนาจการทดสอบของตัวสถิติทดสอบ RML มีค่าสูงที่สุด รองลงมาคือตัวสถิติทดสอบ AD ทุกระดับนัยสำคัญ

2) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K และ KS มีค่าใกล้เคียงกัน

3) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ K มีค่าต่ำสุด สำหรับตัวสถิติทดสอบที่ใช้การแจกแบบลอการิธึม

4) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ KS มีค่าต่ำสุด สำหรับตัวสถิติทดสอบที่ใช้การแจกแบบไวบูลล์

5) ค่าอำนาจการทดสอบของตัวสถิติทดสอบ AD, K, KS และ RML จะสูงขึ้น เมื่อขนาดตัวอย่างเพิ่มขึ้น ทุกระดับนัยสำคัญ

6) อำนาจการทดสอบของตัวสถิติทดสอบทุกตัวสูงขึ้น เมื่อระดับนัยสำคัญเพิ่มขึ้น ในทุกกรณี ซึ่งเป็นผลจากการที่เรากำหนดขอบเขตสมมุติฐานกว้างขึ้น มีผลทำให้ตัวสถิติทดสอบ มีโอกาสที่จะปฏิเสธสมมุติฐาน H_0 ได้มากขึ้น

5.2 ปัญหาที่เกิดขึ้นในการวิจัย

เนื่องจากโปรแกรมที่ใช้เป็นโปรแกรมที่เป็นที่นิยมในสากล แต่ผู้วิจัยไม่มีความชำนาญในการเขียนโปรแกรมจึงต้องใช้เวลาศึกษาและสร้างโปรแกรมเป็นเวลานาน

5.3 ข้อเสนอแนะ

5.3.1 ด้านการเลือกตัวสถิติทดสอบ

จากการวิจัยครั้งนี้ จะพิจารณาตัวสถิติที่สามารถควบคุม ความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1 ได้มากที่สุด และมีอำนาจการทดสอบสูงสุด ดังนั้นควรเลือก สถิติทดสอบ Ratio of Maximum Likelihoods (RML) สำหรับการจำแนกการแจกแบบลอกนอร์มัล และการแจกแจงแบบไวบูลล์

5.3.2 ด้านการวิจัย

5.3.2.1 ควรศึกษาวิจัยตัวสถิติอื่น ๆ ที่ใช้สำหรับการจำแนกการแจกแบบลอกนอร์มัล และการแจกแจงแบบไวบูลล์ เช่น Tiku Test, Watson Test, Cramer von Mises Test ,Shapiro-Wilk Test หรือ Srinivasan Test เพื่อนำมาเปรียบเทียบกับตัวสถิติทดสอบที่ผู้วิจัยได้ศึกษาไว้

5.3.2.2 ควรศึกษาสถิติทดสอบ สำหรับจำแนกการแจกแจงอื่นที่มีความใกล้เคียง เช่น การแจกแจงนอร์มัล (Normal Distribution) และ การแจกแจงลอกนอร์มัล (Lognormal Distribution) เพื่อวิเคราะห์ข้อมูล เมื่อตัวอย่างที่สุ่มมาจากประชากรทราบว่าการแจกแจงอย่างใดอย่างหนึ่ง

บรรณานุกรม

ภาษาไทย

- ประชุม สวัสดิ์. **ทฤษฎีการอนุมานเชิงสถิติ**. กรุงเทพมหานคร : สำนักพิมพ์โอเดียนสโตร์ 2527.
- ศิริรัตน์ วงศ์ประกรณ์กุล. **การทดสอบการแจกแจงไวบูลล์และการแจกแจงกอมเพริตซ์ ด้วยวิธีทดสอบเทียบความกลมกลืนเมื่อข้อมูลถูกตัดทิ้งอย่างมาก**. วิทยานิพนธ์ สถิติศาสตรมหาบัณฑิต ภาควิชาสถิติ บัณฑิตวิทยาลัย จุฬาลงกรณ์มหาวิทยาลัย, 2539.

ภาษาอังกฤษ

- Aslan, B. and Zech, G. **Comparison of Different Goodness-of-Fit Tests**. Durham (ed), Advanced statistical techniques in particle physics. 2002 : 166-175
- Bain, L. J. and Engelhardt, M. **Introduction to Probability and Mathematical Statistics 2nd ed**. New York : Duxbury Press. 1991.
- McDonald, G. C., Vance, L. C. and Gibbons, D.I. **Some Tests for Discriminating Between Lognormal and Weibull Distributions-An Application to Emission Data**. Balakrishnan, N. (ed). Recent Advances in Life-Testing and Reliability. New York. 1995.
- Berr, D. R. and Shudde, R. H. **A note on Kuiper's V_n statistic**. Biometrika. 1973 ; 76 : 663 – 664
- Chernobai, A., Rachev, S. and Fabozzi, F. **Composite goodness-of-fit Tests for left-truncated loss samples**. Technical report, University of California Santa Barbara. 2005.
- Coles, S. G. **On goodness-of-fit tests for the two-parameter Weibull distribution derived from the stabilized probability plot**. Biometrika. 1989 ; 76 : 593 –598
- Crow, E. L. and Shimizu, K. (Eds.). **Lognormal Distributions: Theory and Application**. New York : Marcel-Dekker. 1988.
- D'Agostino, R. B. and Stephens, M. A. **Goodness Of Fit Techniques**. New York : Marcel-Dekker. 1986.

- Dumonceaux, R. and Antle, C. E. **Discrimination between the Log-normal and Weibull distribution.** Technometrics. 1973 ; 15 : 923 – 926.
- _____, Antle, C. E. and Hass, G. **Likelihood ratio test for discrimination between two models with unknown location and scale parameters.** Technometrics. 1973 ; 15 : 923 – 926.
- Larsen, R. I. **Relating air pollutant effects to concentration and control.** Journal of the Air Pollution Control Association. 1970 ; 20 : 214-224.
- Lilliefors, H. W. **On the Kolmogorov-Smirnov test for the exponential distribution with mean unknown.** Journal of the American Statistical Association. 1969 ; 64 : 387-389.
- McDonald, G. C. **Book review of Lognormal Distribution: Theory and Application.** Journal of the American Statistical Association. 1989 ; 84 : 845 – 846.
- Stephens, M. A. **EDF Statistics for Goodness of Fit and Some Comparisons.** Journal of the American Statistical Association. 1976 ; 69 : 730-737.
- Steele, M. C., Chaseling, J., Hurst, C. **A power study of goodness-of-fit tests for categorical data.** International Statistical Institute. 2005 ; 55th Session.
- Shimokawa, T. and Liao, M. **Goodness-of-fit tests for type-I extreme-value and 2-parameter weibull distributions.** IEEE Transactions on Reliability. 1999 ; 48 : 79-86.

ภาคผนวก ก

โปรแกรมที่ใช้ในการสร้างการแจกแจงที่ใช้ในการวิจัย

ในการจำลองชุดตัวอย่างขึ้นมาโดยที่ประชากรมีการแจกแจงตามที่กำหนดไว้โดยอาศัยเทคนิคโปรแกรม R รายละเอียดในการสร้างการแจกแจงแบบต่าง ๆ เป็นดังนี้

การผลิตเลขสุ่มที่มีการแจกแจงแบบไวบูลล์

การแจกแจงแบบไวบูลล์มีฟังก์ชันความหนาแน่นอยู่ในรูปของ

$$g(x;\beta,\alpha) = \frac{\alpha}{\beta} \left(\frac{x}{\beta}\right)^{\alpha-1} \exp\left[-\left(\frac{x}{\beta}\right)^{\alpha}\right], \quad x > 0, \beta > 0, \alpha > 0$$

เมื่อ β เป็น scale parameter

α เป็น shape parameter

การสร้างตัวแปรสุ่มให้มีการแจกแจงไวบูลล์ โดยใช้โปรแกรม R ซึ่งคำสั่งในการสร้าง ตัวแปรให้มีการแจกแจงไวบูลล์ จำนวน n ตัว ที่พารามิเตอร์ $\beta = b$ และ $\alpha = c$ คือ `rweibull (n, c ,b)` ดังตัวอย่างต่อไปนี้

```
> rlnorm
function (n, meanlog = 0, sdlog = 1)
.Internal(rlnorm(n, meanlog, sdlog))
<environment: namespace:stats>
> rweibull(20,0.5,1)
[1] 5.1755583635 0.0366120499 0.2183311546 0.0772656155 0.0001667598
[6] 0.0654551852 0.0045679611 0.0071867150 0.0378979935 0.0225855857
[11] 0.5057724261 0.3996091438 0.8400628308 1.5302615961 0.1395647976
[16] 0.0159737345 8.1828089643 0.2581195054 2.9580998477 1.5356178105
```

การผลิตเลขสุ่มที่มีการแจกแจงแบบล็อกนอร์มัล

การแจกแจงแบบล็อกนอร์มัลมีฟังก์ชันความหนาแน่นอยู่ในรูปของ

$$f(x;\mu,\sigma) = \frac{1}{x\sqrt{2\pi\sigma^2}} \exp\left[-(\ln x - \mu)^2 / 2\sigma^2\right], \quad x > 0, \sigma > 0$$

เมื่อ σ เป็น scale parameter

μ เป็น shape parameter

การสร้างตัวแปรสุ่มให้มีการแจกแจงล็อกนอร์มัล โดยใช้โปรแกรม R ซึ่งคำสั่งในการสร้าง ตัวแปรให้มีการแจกแจงไวบูลล์จำนวน n ตัว ที่พารามิเตอร์ $\sigma = b$ และ $\mu = c$ คือ `rlnorm(n, c, b)` ดังตัวอย่างต่อไปนี้

```
> rweibull
function (n, shape, scale = 1)
.Internal(rweibull(n, shape, scale))
<environment: namespace:stats>
> rlnorm(20,1,0.5)
[1] 1.4042463 2.2420105 3.1535787 3.6548522 4.8233718 1.9393909 5.2188760
[8] 8.0466736 2.5247879 2.7275269 3.5761696 1.2750229 1.0618689 1.8664705
[15] 0.9356022 1.1947580 4.6557036 2.8146571 3.9825933 8.5046730
```

ภาคผนวก ข.

โปรแกรมในการหาค่าความน่าจะเป็นของความคลาดเคลื่อนประเภทที่ 1

Type I error

Lognormal distribution

```
library(MASS)
```

```
library(SuppDists)
```

```
library(statmod)
```

```
adtest.lnorm <- function(x) {
```

```
  x <- sort(x) (เรียงข้อมูล)
```

```
  n <- length(x) (ให้ n คือจำนวนข้อมูล X)
```

```
  result.fitdist <- fitdistr(x, densfun = dlnorm, start=list(meanlog=1, sdlog=1))
```

```
  z <- plnorm(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
```

```
  z.log <- log(z)
```

```
  y <- (2*(1:n) - 1)/n
```

```
  w <- log(1-z[n+1-(1:n)])
```

```
  a <- y * (z.log + w)
```

```
  ad.star <- -sum(a)-n
```

```
  ad.star <- ad.star*(1+(0.75/n)+(2.25/(n^2)))
```

```
  result <- NULL
```

```
  result[1] <- ad.star
```

```
  if(ad.star<0.631) {
```

```
    result[2] = 0
```

```
    result[3] = 0
```

```
    result[4] = 0
```

```
  } else if(ad.star<0.752) {
```

```
    result[2] = 1
```

```
    result[3] = 0
```

```
    result[4] = 0
```

```
  } else if(ad.star<1.035) {
```

```
    result[2] = 1
```

```
    result[3] = 1
```

```
    result[4] = 0
```

```

    } else if(ad.star>=1.035) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

kuiper.Inorm <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun = dlnorm, start=list(meanlog=1, sdlog=1))
    z <- plnorm(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
    D.plus_n_raw = double(n)
    for(i in 1:n) {
        D.plus_n_raw[i] <- (i/n - z[i] )
    }
    D.plus_n <- max(D.plus_n_raw)
    D.minus_n_raw = double(n)
    for(i in 1:n) {
        D.minus_n_raw[i] <- (z[i] - (i-1)/n)
    }
    D.minus_n <- max(D.minus_n_raw)
    k <- D.plus_n + D.minus_n
    result <- NULL
    result[1] <- (sqrt(n)+0.05+(0.82/sqrt(n))) * k

    if(result[1]<1.386) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.489) {
        result[2] = 1
        result[3] = 0
    }

```

```

        result[4] = 0
    } else if(result[1]<1.693) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    } else if(result[1]>=1.693) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

ks.lnorm <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun = dlnorm, start=list(meanlog=1, sdlog=1))
    result <- NULL
    result[1] <- ks.test(x, "plnorm", result.fitdist$estimate[1],
    result.fitdist$estimate[2])$statistic
    result[1] <- result[1]*(sqrt(n)-0.1+(0.85/sqrt(n)))

    if(result[1]<0.819) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<0.895) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.035) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    }
}

```

```

    } else if(result[1]>=1.035) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

rml.lnorm <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun = dweibull, start=list(shape=2, scale=1))
    b <- result.fitdist$estimate[2]
    c <- result.fitdist$estimate[1]
    fw <- c*((x/b)^(c-1))*(1/b) * exp(-((x/b)^c))
    d <- cumprod(x*fw)[n]
    e <- d^(1/n)
    x.ln <- log(x)
    mu.hat <- sum(x.ln)/n
    sigma.hat <- sum((x.ln-mu.hat)^2)/n
    result <- NULL
    result[1] <- (2 * pi * exp(1) * sigma.hat)^(0.5) * e

    if(result[1]<1.03*) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.082*) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.144*) {
        result[2] = 1
        result[3] = 1
    }
}

```

```

        result[4] = 0
    } else if(result[1]>=1.144*) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

```

*เปลี่ยนตามค่า n ตาม ตารางที่ 2-8

ตารางที่ 2-8 ค่า $(RML)_c^{1/n}$ สำหรับ H_0 : การแจกแจงลอการิธึม
และ H_1 : การแจกแจงไวบูลล์

n	$\alpha = 0.10$	$\alpha = 0.05$	$\alpha = 0.01$
20	1.038	1.082	1.144
30	1.020	1.044	1.095
40	1.007	1.028	1.070
50	0.998	1.014	1.054

```

run_ae20 <- function(nlnorm, mu, sigma, n = 1000) {
    result.adtest <- array(0, c(n,4))
    result.kuiper <- array(0, c(n,4))
    result.ks      <- array(0, c(n,4))
    result.rml     <- array(0, c(n,4))

    for(i in 1:n) {
        x <- rlnorm(nlnorm, mu, sigma)
        temp <- adtest.lnorm(x)
        result.adtest[i,] <- c(temp[1], temp[2], temp[3], temp[4])

        temp <- kuiper.lnorm(x)
        result.kuiper[i,] <- c(temp[1], temp[2], temp[3], temp[4])

        temp <- ks.lnorm(x)
        result.ks[i,] <- c(temp[1], temp[2], temp[3], temp[4])
    }
}

```

```

temp <- rml.lnorm(x)
result.rml[i,] <- c(temp[1], temp[2], temp[3], temp[4])
}

adtest <- c(sum(result.adtest[,2])/n, sum(result.adtest[,3])/n,
sum(result.adtest[,4])/n)

kuiper <- c(sum(result.kuiper[,2])/n, sum(result.kuiper[,3])/n,
sum(result.kuiper[,4])/n)

ks <- c(sum(result.ks[,2])/n, sum(result.ks[,3])/n, sum(result.ks[,4])/n)
rml <- c(sum(result.rml[,2])/n, sum(result.rml[,3])/n, sum(result.rml[,4])/n)
}

```

.....
Type I error

Weibull distribution

.....

```

library(MASS)
library(SuppDists)
library(statmod)

```

```

adtest.weibull <- function(x) {
  x <- sort(x)
  n <- length(x)
  result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=2, scale=1))
  z <- pweibull(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
  z.log <- log(z)
  y <- (2*(1:n) - 1)/n
  w <- log(1-z[n+1-(1:n)])
  a <- y * (z.log + w)
  ad.star <- -sum(a)-n
  ad.star <- ad.star*(1+(0.2/sqrt(n)))
  result <- NULL
  result[1] <- ad.star
}

```

```

    if(ad.star<0.637) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(ad.star<0.757) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(ad.star<1.038) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    } else if(ad.star>=1.038) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

kuiper.weibull <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=2, scale=1))
    z <- pweibull(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
    D.plus_n_raw = double(n)
    for(i in 1:n) {
        D.plus_n_raw[i] <- (i/n - z[i] )
    }
    D.plus_n <- max(D.plus_n_raw)
    D.minus_n_raw = double(n)
    for(i in 1:n) {
        D.minus_n_raw[i] <- (z[i] - (i-1)/n)
    }
}

```

```

D.minus_n <- max(D.minus_n_raw)
k <- D.plus_n + D.minus_n
result <- NULL
result[1] <- (sqrt(n)) * k
if(result[1]<1.372) {
  result[2] = 0
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.477) {
  result[2] = 1
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.671) {
  result[2] = 1
  result[3] = 1
  result[4] = 0
} else if(result[1]>=1.671) {
  result[2] = 1
  result[3] = 1
  result[4] = 1
}
return(result)
}

ks.weibull <- function(x) {
  x <- sort(x)
  n <- length(x)
  result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=2, scale=1))
  result <- NULL
  result <- ks.test(x, "pweibull", result.fitdist$estimate[1],
  result.fitdist$estimate[2])$statistic
  result[1] <- result*(sqrt(n))

  if(result[1]<0.803) {

```

```

        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<0.874) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.007) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    } else if(result[1]>=1.007) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

rml.weibull <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=2, scale=1))
    b <- result.fitdist$estimate[2]
    c <- result.fitdist$estimate[1]
    fw <- c*((x/b)^(c-1))*(1/b) * exp(-((x/b)^c))
    d <- cumprod(x*fw)[n]
    e <- d^(1/n)
    x.ln <- log(x)
    mu.hat <- sum(x.ln)/n
    sigma.hat <- sum((x.ln-mu.hat)^2)/n
    result <- NULL
    result[1] <- 1/((2 * pi * exp(1) * sigma.hat)^(0.5) * e)

```

```

if(result[1]<1.041*) {
  result[2] = 0
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.067*) {
  result[2] = 1
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.120*) {
  result[2] = 1
  result[3] = 1
  result[4] = 0
} else if(result[1]>=1.120*) {
  result[2] = 1
  result[3] = 1
  result[4] = 1
}
return(result)
}

```

*เปลี่ยนตามค่า n ตาม ตารางที่ 2-9

ตารางที่ 2-9 ค่า $(RML)_c^{1/n}$ สำหรับ H_0 : การแจกแจงไวบูลล์
และ H_1 : การแจกแจงลอกนอร์มัล

n	$\alpha = 0.10$	$\alpha = 0.05$	$\alpha = 0.01$
20	1.041	1.067	1.120
30	1.019	1.041	1.088
40	1.005	1.026	1.063
50	0.995	1.016	1.045

```

run_weibull20 <- function(nweibull, shape, scale, n = 1000) {
  result.adtest <- array(0, c(n,4))
  result.kuiper <- array(0, c(n,4))
  result.ks      <- array(0, c(n,4))

```

```

result.rml      <- array(0, c(n,4))
for(i in 1:n) {
  x <- rweibull(nweibull, shape, scale)
  temp <- adtest.weibull(x)
  result.adtest[i,] <- c(temp[1], temp[2], temp[3], temp[4])

  temp <- kuiper.weibull(x)
  result.kuiper[i,] <- c(temp[1], temp[2], temp[3], temp[4])
  temp <- ks.weibull(x)
  result.ks[i,] <- c(temp[1], temp[2], temp[3], temp[4])
  temp <- rml.weibull(x)
  result.rml[i,] <- c(temp[1], temp[2], temp[3], temp[4])
}

adtest <- c(sum(result.adtest[,2])/n, sum(result.adtest[,3])/n,
sum(result.adtest[,4])/n)
kuiper <- c(sum(result.kuiper[,2])/n, sum(result.kuiper[,3])/n,
sum(result.kuiper[,4])/n)
ks      <- c(sum(result.ks[,2])/n, sum(result.ks[,3])/n, sum(result.ks[,4])/n)
rml     <- c(sum(result.rml[,2])/n, sum(result.rml[,3])/n, sum(result.rml[,4])/n)

```

ภาคผนวก ค.

โปรแกรมในการหาค่าอำนาจการทดสอบ

.....

```
if (SLIP1 < SLIP2) {
    if (SLIP1 < 1.5) {
        SLIP1 = 1.5;
    }
}
```

$$= 1 - \alpha \log_{10} \left(\frac{\text{result_StdDev} @ \text{location} [51] - \text{result_StdDev} @ \text{location} [50]}{\text{result_StdDev} @ \text{location} [50]} \right)$$
$$x \leq (2*(1-\alpha) - 1)/\alpha$$

3. $\hat{\beta}_1 = 0.0001$ if $\hat{\beta}_1 = 0.0001$ (1.005) $\hat{\beta}_1$

```

    } else if(ad.star>=1.035) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

kuiper.Inorm <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun = dlnorm, start=list(meanlog=1, sdlog=1))
    z <- plnorm(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
    D.plus_n_raw = double(n)
    for(i in 1:n) {
        D.plus_n_raw[i] <- (i/n - z[i] )
    }
    D.plus_n <- max(D.plus_n_raw)
    D.minus_n_raw = double(n)
    for(i in 1:n) {
        D.minus_n_raw[i] <- (z[i] - (i-1)/n)
    }
    D.minus_n <- max(D.minus_n_raw)
    k <- D.plus_n + D.minus_n
    result <- NULL
    result[1] <- (sqrt(n)+0.05+(0.82/sqrt(n))) * k

    if(result[1]<1.386) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.489) {
        result[2] = 1
        result[3] = 0
    }

```

```

        result[4] = 0
    } else if(result[1]<1.693) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    } else if(result[1]>=1.693) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

ks.lnorm <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun = dlnorm, start=list(meanlog=1, sdlog=1))
    result <- NULL
    result[1] <- ks.test(x, "plnorm", result.fitdist$estimate[1],
    result.fitdist$estimate[2])$statistic
    result[1] <- result[1]*(sqrt(n)-0.1+(0.85/sqrt(n)))

    if(result[1]<0.819) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<0.895) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.035) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    }
}

```

```

    } else if(result[1]>=1.035) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

rml.lnorm <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun = dweibull, start=list(shape=2, scale=1))
    b <- result.fitdist$estimate[2]
    c <- result.fitdist$estimate[1]
    fw <- c*((x/b)^(c-1))*(1/b) * exp(-((x/b)^c))
    d <- cumprod(x*fw)[n]
    e <- d^(1/n)
    x.ln <- log(x)
    mu.hat <- sum(x.ln)/n
    sigma.hat <- sum((x.ln-mu.hat)^2)/n
    result <- NULL
    result[1] <- (2 * pi * exp(1) * sigma.hat)^(0.5) * e

    if(result[1]<1.038*) {
        result[2] = 0
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.082*) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.144*) {
        result[2] = 1
        result[3] = 1
    }
}

```

```

        result[4] = 0
    } else if(result[1]>=1.144*) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

```

*เปลี่ยนตามค่า n ตาม ตารางที่ 2-8

ตารางที่ 2-8 ค่า $(RML)_c^{1/n}$ สำหรับ H_0 : การแจกแจงลอกนอร์มัล
และ H_1 : การแจกแจงไวบูลล์

n	$\alpha = 0.10$	$\alpha = 0.05$	$\alpha = 0.01$
20	1.038	1.082	1.144
30	1.020	1.044	1.095
40	1.007	1.028	1.070
50	0.998	1.014	1.054

```

run_power120 <- function(nweibull, shape, scale, n = 1000) {
    result.adtest <- array(0, c(n,4))
    result.kuiper <- array(0, c(n,4))
    result.ks <- array(0, c(n,4))
    result.rml <- array(0, c(n,4))
    for(i in 1:n) {
        x <- rweibull(nweibull, shape, scale)
        temp <- adtest.lnorm(x)
        result.adtest[i,] <- c(temp[1], temp[2], temp[3], temp[4])
        temp <- kuiper.lnorm(x)
        result.kuiper[i,] <- c(temp[1], temp[2], temp[3], temp[4])
        temp <- ks.lnorm(x)
        result.ks[i,] <- c(temp[1], temp[2], temp[3], temp[4])
        temp <- rml.lnorm(x)
        result.rml[i,] <- c(temp[1], temp[2], temp[3], temp[4])
    }
}

```

```

}

adtest <- c(sum(result.adtest[,2])/n, sum(result.adtest[,3])/n,
sum(result.adtest[,4])/n)
kuiper <- c(sum(result.kuiper[,2])/n, sum(result.kuiper[,3])/n,
sum(result.kuiper[,4])/n)
ks <- c(sum(result.ks[,2])/n, sum(result.ks[,3])/n, sum(result.ks[,4])/n)
rml <- c(sum(result.rml[,2])/n, sum(result.rml[,3])/n, sum(result.rml[,4])/n)
}

```

.....
Power of the test

Weibull distribution
.....

```

library(MASS)
library(SuppDists)
library(statmod)
adtest.weibull <- function(x) {
  x <- sort(x)
  n <- length(x)
  result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=1, scale=1))
  z <- pweibull(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
  z.log <- log(z)
  y <- (2*(1:n) - 1)/n
  w <- log(1-z[n+1-(1:n)])
  a <- y * (z.log + w)
  ad.star <- -sum(a)-n
  ad.star <- ad.star*(1+(0.2/sqrt(n)))
  result <- NULL
  result[1] <- ad.star
  if(ad.star<0.637) {
    result[2] = 0
  }
}

```

```

        result[3] = 0
        result[4] = 0

    } else if(ad.star<0.757) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(ad.star<1.038) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    } else if(ad.star>=1.038) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }
    return(result)
}

kuiper.weibull <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=1, scale=1))
    z <- pweibull(x, result.fitdist$estimate[1], result.fitdist$estimate[2])
    D.plus_n_raw = double(n)
    for(i in 1:n) {
        D.plus_n_raw[i] <- (i/n - z[i] )
    }
    D.plus_n <- max(D.plus_n_raw)
    D.minus_n_raw = double(n)
    for(i in 1:n) {
        D.minus_n_raw[i] <- (z[i] - (i-1)/n)
    }
    D.minus_n <- max(D.minus_n_raw)

```

```

k <- D.plus_n + D.minus_n
result <- NULL
result[1] <- (sqrt(n)) * k

if(result[1]<1.372) {
  result[2] = 0
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.477) {
  result[2] = 1
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.671) {
  result[2] = 1
  result[3] = 1
  result[4] = 0
} else if(result[1]>=1.671) {
  result[2] = 1
  result[3] = 1
  result[4] = 1
}
return(result)
}

ks.weibull <- function(x) {
  x <- sort(x)
  n <- length(x)
  result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=1, scale=1))
  result <- NULL
  result <- ks.test(x, "pweibull", result.fitdist$estimate[1],
  result.fitdist$estimate[2])$statistic
  result[1] <- result*(sqrt(n))

  if(result[1]<0.803) {

```

```

        result[2] = 0
        result[3] = 0
        result[4] = 0

    } else if(result[1]<0.874) {
        result[2] = 1
        result[3] = 0
        result[4] = 0
    } else if(result[1]<1.007) {
        result[2] = 1
        result[3] = 1
        result[4] = 0
    } else if(result[1]>=1.007) {
        result[2] = 1
        result[3] = 1
        result[4] = 1
    }

    return(result)
}

rml.weibull <- function(x) {
    x <- sort(x)
    n <- length(x)
    result.fitdist <- fitdistr(x, densfun=dweibull, start=list(shape=1, scale=1))
    b <- result.fitdist$estimate[2]
    c <- result.fitdist$estimate[1]
    fw <- c*((x/b)^(c-1))*(1/b) * exp(-((x/b)^c))
    d <- cumprod(x*fw)[n]
    e <- d^(1/n)
    x.ln <- log(x)
    mu.hat <- sum(x.ln)/n
    sigma.hat <- sum((x.ln-mu.hat)^2)/n
    result <- NULL

```

```
result[1] <- 1/((2 * pi * exp(1) * sigma.hat)^(0.5) * e)
```

```
if(result[1]<1.041*) {
  result[2] = 0
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.067*) {
  result[2] = 1
  result[3] = 0
  result[4] = 0
} else if(result[1]<1.120*) {
  result[2] = 1
  result[3] = 1
  result[4] = 0
} else if(result[1]>=1.120*) {
  result[2] = 1
  result[3] = 1
  result[4] = 1
}
return(result)
```

```
}
```

*เปลี่ยนตามค่า n ตาม ตารางที่ 2-9

ตารางที่ 2-9 ค่า $(RML)_c^{1/n}$ สำหรับ H_0 : การแจกแจงไวบูลล์
และ H_I : การแจกแจงลอกนอร์มัล

n	$\alpha = 0.10$	$\alpha = 0.05$	$\alpha = 0.01$
20	1.041	1.067	1.120
30	1.019	1.041	1.088
40	1.005	1.026	1.063
50	0.995	1.016	1.045

```
run_powerw20 <- function(nlnorm, mu, sigma, n = 1000) {
```

```

result.adtest <- array(0, c(n,4))
result.kuiper <- array(0, c(n,4))
result.ks <- array(0, c(n,4))
result.rml <- array(0, c(n,4))

for(i in 1:n) {
  x <- rlnorm(nlnorm, mu, sigma)
  temp <- adtest.weibull(x)
  result.adtest[i,] <- c(temp[1], temp[2], temp[3], temp[4])

  temp <- kuiper.weibull(x)
  result.kuiper[i,] <- c(temp[1], temp[2], temp[3], temp[4])
  temp <- ks.weibull(x)
  result.ks[i,] <- c(temp[1], temp[2], temp[3], temp[4])
  temp <- rml.weibull(x)
  result.rml[i,] <- c(temp[1], temp[2], temp[3], temp[4])
}

adtest <- c(sum(result.adtest[,2])/n, sum(result.adtest[,3])/n,
sum(result.adtest[,4])/n)

kuiper <- c(sum(result.kuiper[,2])/n, sum(result.kuiper[,3])/n,
sum(result.kuiper[,4])/n)

ks <- c(sum(result.ks[,2])/n, sum(result.ks[,3])/n, sum(result.ks[,4])/n)
rml <- c(sum(result.rml[,2])/n, sum(result.rml[,3])/n, sum(result.rml[,4])/n)
}

```

ประวัติผู้วิจัย

ชื่อ : นางสาวศรีอัมพร เร่บ้านเกาะ
 ชื่อวิทยานิพนธ์ : การเปรียบเทียบอำนาจการทดสอบภาวะสารูปดีของสถิติบางตัว
 สำหรับการแจกแจงลอการิธึมและไวบูลล์
 สาขาวิชา : สถิติประยุกต์

ประวัติ

ประวัติส่วนตัวเกิดวันที่ 4 ธันวาคม พ.ศ. 2523

ประวัติการศึกษา สำเร็จการศึกษาระดับปริญญาตรี การศึกษามหาบัณฑิต(กศ.บ.) จากมหาวิทยาลัยศรีนครินทรวิโรฒ คณะศึกษาศาสตร์ เอกคณิตศาสตร์ ในปี 2544 และเข้าศึกษาในหลักสูตรวิทยาศาสตรมหาบัณฑิต ในสถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ คณะวิทยาศาสตร์ประยุกต์ ภาควิชาสถิติประยุกต์ ในปีการศึกษา 2546