

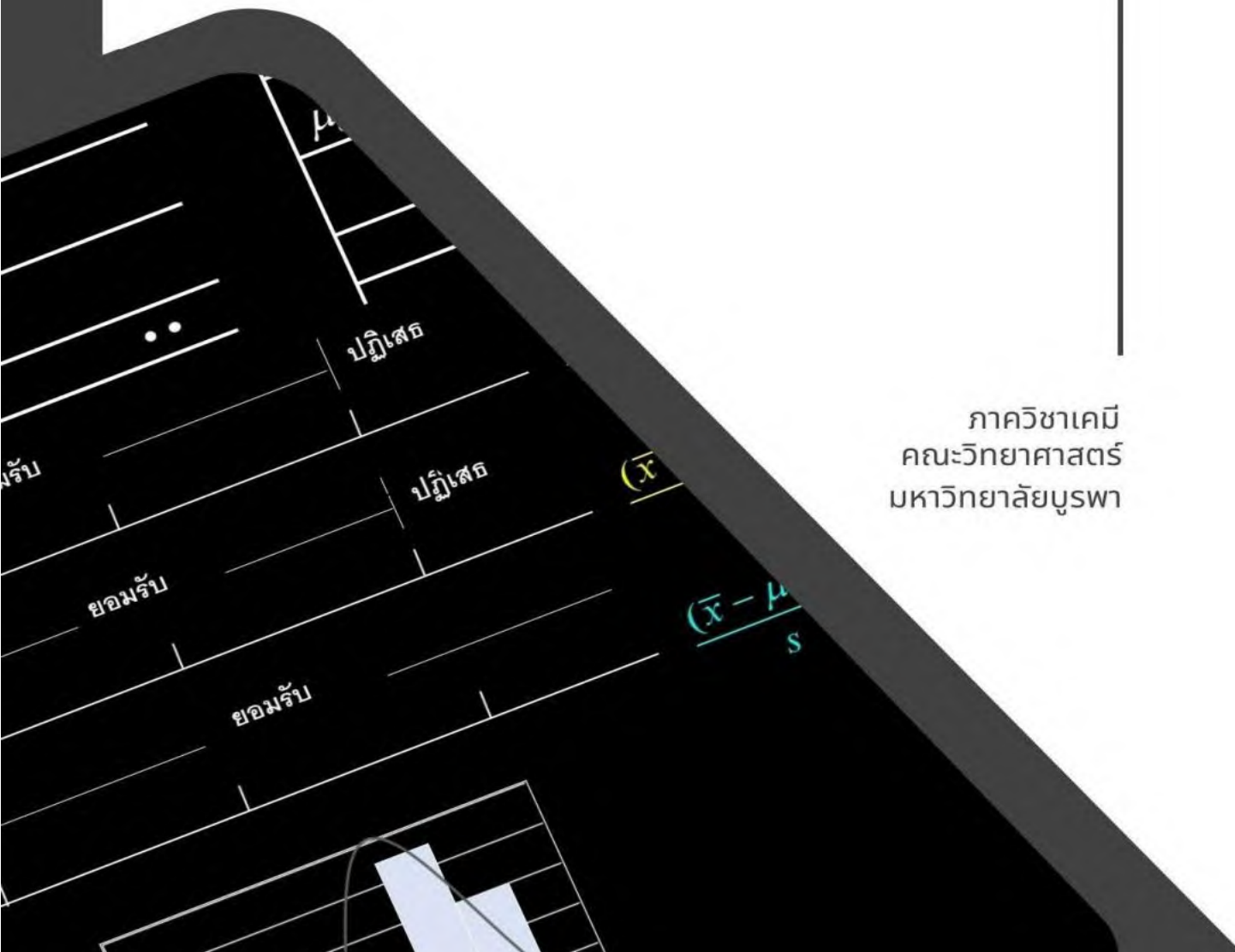
ผู้ช่วยศาสตราจารย์ ดร. นภา ตั้งเตริยมจิตมั่น

สถิติสำหรับเคมีวิเคราะห์

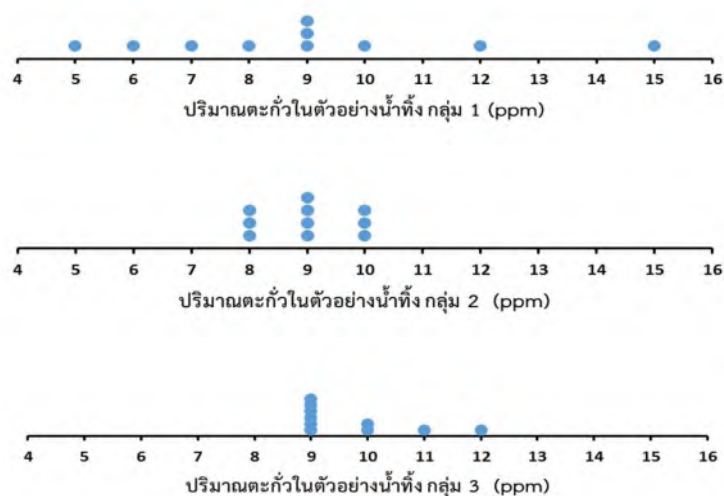
STATISTICS FOR

Analytical Chemistry

ภาควิชาเคมี
คณะวิทยาศาสตร์
มหาวิทยาลัยบูรพา



ตำราวิชา 30353464
สถิติสำหรับเคมีวิเคราะห์
Statistics for Analytical Chemistry



*Are they significantly different?
Statistics can tell you !!!*

ผู้ช่วยศาสตราจารย์ ดร. นภา ตั้งเตรียมจิตมั่น

ภาควิชาเคมี คณะวิทยาศาสตร์
มหาวิทยาลัยบูรพา จ.ชลบุรี
2567

คำนำ

ตำราเล่มนี้จัดทำขึ้นเพื่อใช้สอนวิชา 30353464 สถิติสำหรับเคมีวิเคราะห์ (Statistics for Analytical Chemistry) ซึ่งเป็นวิชาในหลักสูตรวิทยาศาสตรมหาบัณฑิต สาขาวิชาเคมี ฉบับปรับปรุง พ.ศ. 2564 หมวดวิชาเลือก กลุ่มวิชาเคมีวิเคราะห์ สำหรับนิสิตภาควิชาเคมี คณะวิทยาศาสตร์ ระดับบัณฑิตศึกษา

งานทางเคมีวิเคราะห์จะมุ่งเน้นการวิเคราะห์หาปริมาณสาร การพัฒนาวิธีวิเคราะห์และการตรวจสอบว่าวิธีวิเคราะห์ที่พัฒนาขึ้นมีความถูกต้อง น่าเชื่อถือ เป็นงานหลักของนักวิจัยและพัฒนา ส่วนผู้ที่ทำงานวิเคราะห์ประจำวันจะมีการทวนสอบประสิทธิภาพของวิธีซึ่งเป็นหนึ่งในกิจกรรมการประกันคุณภาพทางเคมีวิเคราะห์ ขั้นตอนเหล่านี้จะเริ่มตั้งแต่การหาสภาวะที่เหมาะสมในการเตรียมตัวอย่าง การตรวจวัด และจบด้วยการตรวจสอบประสิทธิภาพของวิธี ล้วนต้องใช้สถิติในการออกแบบการทดลอง เปรียบเทียบหรือตัดสินใจ เพื่อสรุปข้อมูลที่ได้จากการทดลองซึ่งแตกต่างจากข้อมูลที่ได้จากการสำรวจด้วยแบบสอบถาม หรือการเก็บข้อมูลด้วยการสังเกต ข้อมูลทางเคมีวิเคราะห์ จะเกี่ยวข้องกับทั้งผู้ทำการทดลอง ตัวอย่าง สารเคมี อุปกรณ์และเครื่องมือที่ใช้ วัน เวลาและสถานที่ทำการทดลอง เป็นต้น การนำสถิติมาใช้จึงต้องคำนึงถึงปัจจัยเหล่านี้เป็นสำคัญ เพิ่มเติมจากการเข้าใจหลักของวิธีทางสถิติเพื่อนำมาประยุกต์ใช้กับข้อมูลที่ได้จากการทดลอง

เนื้อหาของตำราจะเรียงลำดับบทตามขั้นตอนของการหาปริมาณวิเคราะห์เพื่อหาสารที่สนใจในตัวอย่าง โดยบทที่ 1-3 เป็นการทบทวนความรู้เฉพาะวิธีทางสถิติที่จะนำมาใช้ มีการแสดงการคำนวณเพื่อให้ทราบที่มาของสูตรที่ใช้โปรแกรมเอ็กเซลในการคำนวณของบทที่ 4-7 ซึ่งมีการใช้ข้อมูลจริงจากบทความวิจัยที่ตีพิมพ์ในวารสารมาเป็นตัวอย่างการคำนวณ เป็นการเพิ่มความเข้าใจในการใช้ความรู้ทางสถิติกับข้อมูลทางเคมีวิเคราะห์ อนึ่ง สำหรับตารางค่าวิกฤตจะมีแทรกเป็นฉบับย่อที่เพียงพอต่อการใช้งานตามเนื้อหาที่กล่าวถึง

ผู้เขียนหวังเป็นอย่างยิ่งว่าตำราเล่มนี้จะเป็นประโยชน์ต่อการเรียนการสอนตามสมควร หากนิสิตหรือท่านที่นำไปใช้มีข้อเสนอแนะหรือคำติชม โปรดติดต่อผู้เขียนได้ที่ อีเมล napa@go.buu.ac.th ผู้เขียนยินดีรับฟังข้อคิดเห็นและขอขอบคุณล่วงหน้ามา ณ โอกาสนี้

นภา ตั้งเตรียมจิตมั่น

(ผู้ช่วยศาสตราจารย์ ดร. นภา ตั้งเตรียมจิตมั่น)

7 พฤษภาคม 2567

สารบัญ

บทที่	เรื่อง	หน้า
1	ความสำคัญของสถิติในงานทางเคมีวิเคราะห์	1
	1.1 สถิติเชิงพรรณนา (descriptive statistics)	1
	1.2 การแจกแจงแบบปกติ (normal distribution)	6
	1.3 สถิติเชิงอนุมาน (inferential or inductive statistics)	8
	1.4 การใช้สถิติในงานทางเคมีวิเคราะห์	10
	1.5 ความคลาดเคลื่อนในการทำปริมาณวิเคราะห์	11
	แบบฝึกหัดท้ายบทที่ 1	13
	เอกสารอ้างอิง	14
2	ข้อมูลและการตรวจวัดซ้ำ	15
	2.1 ประเภทของข้อมูล (data type)	18
	2.2 การเก็บตัวอย่างแบบสุ่ม	19
	2.3 การแจกแจงของค่าเฉลี่ยของตัวอย่างที่ได้จากการสุ่ม (distribution of sample mean)	21
	2.4 การตรวจวัดซ้ำ (repeated measurements)	21
	2.5 ช่วงความเชื่อมั่น (confidence interval)	22
	2.6 การคำนวณเกี่ยวกับ ช่วงความเชื่อมั่น	24
	2.7 การคำนวณความคลาดเคลื่อนส่งต่อ (propagation error)	27
	แบบฝึกหัดท้ายบทที่ 2	31
	เอกสารอ้างอิง	32

3	การทดสอบสมมติฐาน	33
	3.1 การทดสอบค่าสุดต่างด้วย Grubbs' test	33
	3.2 การทดสอบค่าความแปรปรวนที่เป็นค่าสุดต่างด้วย Cochran test	36
	3.3 การทดสอบค่าความแปรปรวนด้วยสถิติทดสอบเอฟ (F-test)	38
	3.4 การวิเคราะห์ความแปรปรวนทางเดียว (one-way ANOVA)	43
	3.5 การวิเคราะห์ความแปรปรวนสองทาง (two-way ANOVA)	48
	3.6 การทดสอบค่าเฉลี่ยด้วยสถิติทดสอบที (t-test)	62
	3.7 สถิติทดสอบทีเพื่อเปรียบเทียบประชากร 1 กลุ่ม กับค่าจริง	66
	3.8 สถิติทดสอบทีเพื่อเปรียบเทียบประชากร 2 กลุ่ม	70
	3.9 สถิติทดสอบทีแบบ paired t-test	75
	3.10 การพิสูจน์ $F = t^2$	78
	แบบฝึกหัดท้ายบทที่ 3	80
	เอกสารอ้างอิง	82
4	การหาสถานะที่เหมาะสม	83
	4.1 การออกแบบการทดลองแบบสุ่ม (randomized design)	83
	4.2 การออกแบบการทดลองแบบสุ่มในบล็อก (randomized block design)	88
	4.3 การออกแบบการทดลองแบบจัตุรัสลาติน (Latin square design)	90
	4.4 การออกแบบการทดลองแบบ crossed design	92
	แบบฝึกหัดท้ายบทที่ 4	99
	เอกสารอ้างอิง	100

5	สถิติในวิธีเทียบมาตรฐาน	101
	5.1 การสร้างกราฟมาตรฐาน (calibration curve)	101
	5.2 สัมประสิทธิ์สหสัมพันธ์ (correlation coefficient, r)	102
	5.3 สมการเส้นตรงของกราฟมาตรฐาน	105
	5.4 การทำปริมาณวิเคราะห์ด้วยวิธี standard addition	110
	5.5 การแปลงรูป (data transformation) ข้อมูลเพื่อให้ได้ ความสัมพันธ์เป็นสมการเส้นตรง	116
	5.6 กราฟมาตรฐานถดถอยเชิงเส้นแบบถ่วงน้ำหนัก (weighted regression lines)	117
	แบบฝึกหัดท้ายบทที่ 5	125
	เอกสารอ้างอิง	126

6	ความเที่ยง ความถูกต้อง และขีดจำกัดการตรวจวัด	127
6.1	ความเที่ยง	127
6.2	การหาค่า repeatability และ between-parameter reproducibility พร้อมกัน	128
6.3	หาทดสอบค่าสุดต่างของข้อมูลก่อนการหาความเที่ยง	131
6.4	การหาความเที่ยงประเภท intermediate precision	133
6.5	ขีดจำกัดความเที่ยง (precision limit)	138
6.6	การเปรียบเทียบค่าความเที่ยงของวิธีวิเคราะห์สองวิธี	140
6.7	ความถูกต้อง (accuracy)	143
6.8	การตรวจสอบความถูกต้องของวิธีด้วยสารมาตรฐานอ้างอิง	145
6.9	การตรวจสอบความถูกต้องของวิธีด้วยการเทียบกับวิธีมาตรฐาน	146
6.10	การเปรียบเทียบวิธีวิเคราะห์มากกว่า 2 วิธี ด้วยสถิติวิเคราะห์ความแปรปรวน (ANOVA)	154
6.11	ขีดจำกัดการตรวจวัด	158
6.12	ที่มาทางสถิติของค่า LOD และ LOQ	158
6.13	การหาค่า LOD จากการตรวจวัดสารละลายแบบลงค์	161
6.14	การหาค่า LOD โดยการคำนวณจากกราฟมาตรฐาน	163
	แบบฝึกหัดท้ายบทที่ 6	165
	เอกสารอ้างอิง	169

7	สถิติในงานประกันคุณภาพและควบคุมคุณภาพ	171
	7.1 ยุทธวิธีการสุ่มตัวอย่าง (sampling strategy)	171
	7.2 การหาความเป็นเนื้อเดียวของวัสดุอ้างอิงรับรอง	175
	7.3 การวางแผนชักตัวอย่างเพื่อการยอมรับ (acceptance sampling plan)	178
	7.4 การทดสอบความทน (robustness test)	179
	7.5 การทำ collaborative trial (CT)	185
	7.6 การทดสอบความชำนาญ (proficiency testing; PT)	192
	7.7 แผนภูมิควบคุม (control chart)	195
	แบบฝึกหัดท้ายบทที่ 7	206
	เอกสารอ้างอิง	209
	คำตอบแบบฝึกหัดท้ายบทบางข้อ	211
	บรรณานุกรม	215
	INDEX	218

ความสำคัญของสถิติในงานทางเคมีวิเคราะห์

การศึกษาหรือการวิจัยทางสังคม วิทยาศาสตร์ สิ่งแวดล้อม โภชนาการ อุตสาหกรรม สุขภาพ และกิจกรรมทั่วไปในปัจจุบันนี้ล้วนแต่ได้ข้อมูลที่ได้จากการสำรวจ หรือการทำการทดลอง ซึ่งต้องมีการแปลผลด้วยวิธีทางสถิติ สถิติใช้คณิตศาสตร์เป็นเครื่องมือในการคำนวณ ด้วยการรวบรวมข้อมูล หรือข้อเท็จจริงเชิงตัวเลข มาวิเคราะห์แสดงเป็นข้อสรุปที่ถูกต้อง หรือให้ความหมายเข้าใจได้ง่ายขึ้น เช่น ประชากร 1 ใน 6 ของประเทศไทย อาศัยอยู่ในกรุงเทพ นิสิตที่จบคณะวิทยาศาสตร์ 90% ทำงานในโรงงานอุตสาหกรรม เป็นต้น สถิติใช้วิธีคำนวณทางคณิตศาสตร์อย่างง่าย ตั้งแต่ บวก ลบ คูณ หาร จนถึงการคำนวณอย่างยาก อย่างทฤษฎีความน่าจะเป็นเพื่อวิเคราะห์ข้อมูล สถิติจึงเป็นเครื่องมือช่วยในการตัดสินใจในการศึกษาหรือการทำการวิจัยในเรื่องมากมาย สถิติแบ่งออกเป็น สถิติเชิงพรรณนาและสถิติเชิงอนุมาน ขั้นตอนทางสถิติในการจัดการข้อมูลเริ่มตั้งแต่ การรวบรวมข้อมูล (gathering data) การวิเคราะห์ข้อมูล (analyzing data) การสรุปข้อมูล (summarizing data) และการรายงานผลการวิเคราะห์ (reporting the results) ทุกขั้นตอนใช้สถิติเชิงพรรณนา ยกเว้นการวิเคราะห์ข้อมูลที่ใช้สถิติเชิงอนุมาน ซึ่งใช้ข้อมูลจากกลุ่มตัวอย่างมาอนุมานหรือคาดการณ์เพื่อบ่งบอกลักษณะเฉพาะของประชากร หรือเปรียบเทียบประชากรตั้งแต่ 2 กลุ่มขึ้นไป

1.1 สถิติเชิงพรรณนา (descriptive statistics)

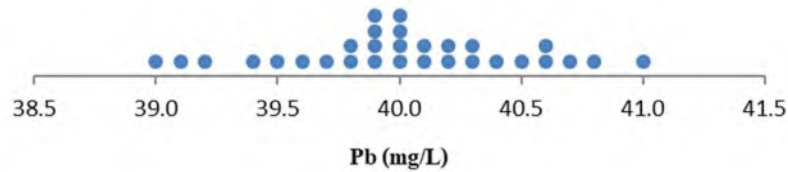
คือสถิติที่ใช้ในการนำเสนอข้อมูล และการรายงานผลการวิเคราะห์ ที่ได้จากการวิเคราะห์กลุ่มตัวอย่าง โดยมีทั้งแบบที่เป็นกราฟิก และแบบที่เป็นค่าตัวเลข สถิติเชิงพรรณนาช่วยให้เห็นภาพรวมหรือข้อสรุปของข้อมูลจากกลุ่มตัวอย่างได้ชัดเจนขึ้น และนำข้อมูลนั้นไปใช้สถิติเชิงอนุมานเพื่อบอกลักษณะเฉพาะของประชากร โดยปกติจะมีข้อแนะนำว่าให้ดูภาพรวมของข้อมูลในรูปแบบกราฟิกเป็นการตัดสินใจเบื้องต้น เพื่อเลือกวิธีทางสถิติที่เหมาะสมในการวิเคราะห์ข้อมูลต่อไป หรือก่อนที่จะลงมือคำนวณค่าสถิติที่เป็นตัวเลข เพราะกราฟิกจะทำให้เห็นแนวโน้ม และสรุปภาพรวม แสดงลักษณะของข้อมูลได้ดีกว่าตัวเลขจำนวนมากที่อยู่รวมกัน กราฟิกจึงเป็นส่วนหนึ่งของการวิเคราะห์ข้อมูลเชิงสำรวจ หรือการวิเคราะห์ข้อมูลเบื้องต้น (exploratory data analysis; EDA หรือ initial data analysis; IDA) กราฟิกที่สร้างง่ายซึ่งจัดเป็น EDA ได้แก่ แผนภาพจุด และแผนภาพลำต้นและใบ

1.1.1 แบบกราฟิก

ใช้วิธีนำเสนอข้อมูลตัวเลขแบบเป็นกราฟ หลากหลายประเภทหรือเป็นแผนภาพ

ก. แผนภาพจุด (dot plot) ทำให้เห็นตำแหน่งสัมพัทธ์ (relative position) ของข้อมูล เหมาะสำหรับข้อมูลจำนวนน้อย เป็นการแสดงข้อมูลแบบต่อเนื่อง แต่ละค่าของข้อมูลที่เท่ากันจะแสดงเป็นจุดวางเรียงกันในแนวตั้งดังตัวอย่างแผนภาพจุด ในรูปที่ 1.1 ของข้อมูลปริมาณตะกั่วในน้ำทิ้ง 30 ค่า ดังต่อไปนี้

39.0	39.1	39.2	39.4	39.5	39.6	39.7	39.8	39.8	39.9
39.9	39.9	39.9	40.0	40.0	40.0	40.0	40.1	40.1	40.2
40.2	40.3	40.3	40.4	40.5	40.6	40.6	40.7	40.8	41.0



รูปที่ 1.1 แผนภาพจุด (dot plot) สร้างด้วยโปรแกรมเอ็กเซล

ข. แผนภาพลำต้นและใบ (stem and leaf diagram) ใช้ตัวเลขแสดงเป็นกราฟ หลักที่มีค่าสูงสุดเป็นแกนใน แนวคอลัมน์ เรียกว่า ลำต้น และหลักของตัวเลขถัดไปจะวางด้านขวาของหลักลำต้น เรียกว่า ใบ ชิดเส้นคั่นระหว่างลำต้นและใบ แผนภาพนี้สร้างง่ายและทำให้เห็นภาพเหมือนของฮิสโทแกรมที่สร้างยากกว่า

Leaf unit: 0.5

39		0 1 2 4
39		5 6 7 8 8 9 9 9 9
40		0 0 0 0 1 1 2 2 3 3 4
40		5 6 6 7 8
41		0

รูปที่ 1.2 แผนภาพลำต้นและใบ (stem and leaf diagram) ของข้อมูลในรูปที่ 1.1

ค. ฮิสโทแกรม (histogram) เป็นกราฟที่ใช้แสดงความถี่ของข้อมูลที่มีค่าแบบต่อเนื่อง แกนนอนแสดงช่วงของข้อมูลและแกนตั้งแสดงค่าของความถี่ของข้อมูลในแต่ละช่วง ฮิสโทแกรมเป็นสถิติแบบกราฟิกที่ช่วยให้เห็นการแจกแจงของข้อมูล โดยการสร้างฮิสโทแกรมต้องคำนึงถึงความกว้างอันตรภาคชั้น (bin width; h) ที่เหมาะสม ซึ่งสามารถคำนวณได้จากสมการที่ 1.1 1.2 หรือ 1.3 และโดยทั่วไปมักนิยมให้มีช่วงอันตรภาคชั้นหรือจำนวนชั้นประมาณ 5-20 ชั้นหรือประมาณ $\sqrt{n}$ (<https://en.wikipedia.org/wiki/Histogram>)

สมการแรกนี้ใช้ค่าสูงสุด (max) ลบด้วยค่าต่ำสุด (min) และนำมาหารด้วยรากที่สองของจำนวนข้อมูล (n)

$$h = \frac{\max - \min}{\sqrt{n}} \quad \dots\dots\dots 1.1$$

โดย h คือ bin width

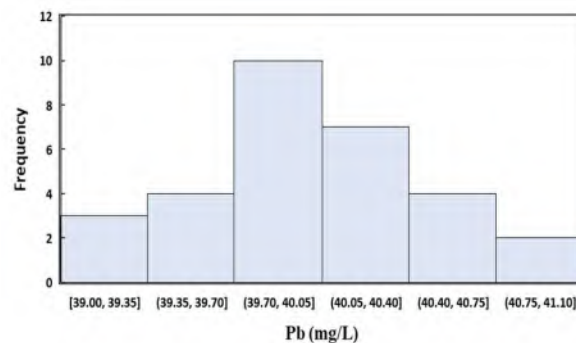
สถิติสำหรับเคมีวิเคราะห์

ถ้าข้อมูลมีมากและมีแนวโน้มว่าจะมีการแจกแจงปกติ $h = \frac{3.5s}{\sqrt[3]{n}}$ 1.2

โดย s คือ ส่วนเบี่ยงเบนมาตรฐาน

ถ้าข้อมูลมีน้อยและมีแนวโน้มว่าจะไม่มีการแจกแจงปกติ $h = \frac{2(IQR)}{\sqrt[3]{n}}$ 1.3

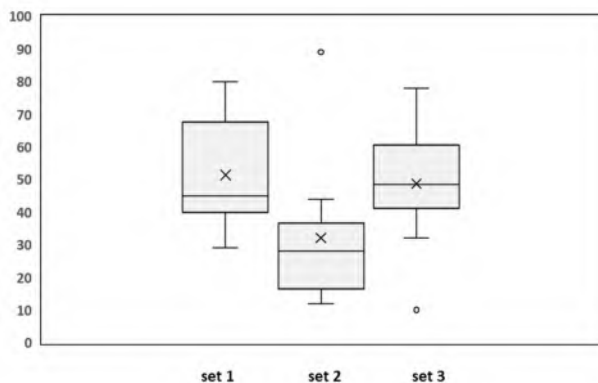
โดยที่ IQR คือ (interquartile range) หรือ ค่าพิสัยควอไทล์ = $Q3 - Q1$



รูปที่ 1.3 แผนภาพฮิสโทแกรม (histogram) ของข้อมูลในรูปที่ 1.1 มี bin width = 0.35 mg/L (สมการ 1.3)

เมื่อนำข้อมูลในรูปที่ 1.1 มาหา bin width ด้วยสมการที่ 1.1 และ 1.3 พบว่าได้ใกล้เคียงกัน และได้จำนวนชั้นเท่ากับ 6 ใกล้เคียงค่า $\sqrt{30}$ ถ้าใช้สมการที่ 1.2 ได้จำนวนชั้นเพียง 4 ชั้น เป็นเพราะจำนวนข้อมูลมีน้อยเกินไปสำหรับการสร้างฮิสโทแกรม ควรมีข้อมูลอย่างน้อย 50 ค่าขึ้นไป ถ้าข้อมูลจำนวนน้อยแผนภาพจะแสดงภาพรวมได้ดีกว่า (Thompson, 2011)

ง. แผนภาพกล่อง (box plot หรือ box-and-whisker plot) เป็นการนำค่าตัวเลขทางสถิติ 5 ค่า มาพลอตเรียงเป็นรูปกราฟจากล่างขึ้นบนตามแกน y ดังนี้ คือ ค่าต่ำสุด ค่าที่ควอไทล์ที่ 1 ($Q1$) ค่าที่ควอไทล์ที่ 2 ($Q2$) ค่าที่ควอไทล์ที่ 3 ($Q3$) และค่าสูงสุด ดังในรูปที่ 1.4 กราฟ box plot จะช่วยให้เห็น การแจกแจงของข้อมูล ข้อมูลชุดที่ 2 มีค่าน้อยกว่าชุดที่ 1 และ 3 นอกจากนี้ยังแสดงค่าสุดต่าง (outlier) ซึ่งเป็นข้อมูลที่ต่างไปจากค่าอื่น ดังเช่นชุดที่ 2 มีค่าสุดต่างที่มีค่ามากกว่าค่าอื่นพลอตอยู่ข้างบน ส่วนชุดที่ 3 มีค่าสุดต่างที่มีค่าน้อยกว่าค่าอื่นพลอตอยู่ข้างล่าง ข้อมูลชุดที่ 3 มีค่าเฉลี่ย (x) เท่ากับค่าที่ควอไทล์ที่ 2 (ค่ามัธยฐาน)



รูปที่ 1.4 แผนภาพกล่อง ของข้อมูล 3 ชุด (set 1, set 2 และ set 3) โดย x คือค่าเฉลี่ย และ ° คือ ค่าสุดต่าง

จ. cumulative distribution เป็นกราฟที่พลอตระหว่างสัดส่วนความถี่ที่พบของข้อมูลและค่าข้อมูลนั้น โดยให้สัดส่วนความถี่ที่พบเป็นแกน Y และค่าข้อมูลเป็นแกน X สัดส่วนความถี่อาจใช้ในรูปแบบเปอร์เซ็นต์ความถี่ก็ได้ ถ้าการแจกแจงของข้อมูลเป็นแบบปกติ กราฟที่ได้จะมีรูปเป็นตัวเอส เรียกว่า cumulative frequency curve

ฉ. normality probability plot เป็นเทคนิคที่ใช้กราฟทดสอบการแจกแจงข้อมูล ว่าข้อมูลมีการแจกแจงแบบปกติ หรือเรียกอีกอย่างว่า การแจกแจงแบบเกาส์เซียน (gaussian distribution) หรือไม่ โดยการพลอต cumulative frequency curve บนกระดาษที่มีแกน Y เป็นลอการิทึมสเกล (normality probability paper) กราฟที่ได้จะออกมาเป็นกราฟเส้นตรง ถ้าเป็นการแจกแจงแบบเกาส์เซียน เนื่องจากสถิติเชิงอนุมานส่วนใหญ่ที่ใช้กันจะเป็นสถิติแบบที่ใช้พารามิเตอร์ ซึ่งมีเงื่อนไขว่าข้อมูลที่น่ามาทดสอบต้องมีการแจกแจงแบบปกติ

1.1.2 แบบค่าตัวเลข

เป็นสถิติเชิงพรรณนามีอยู่ด้วยกัน 3 กลุ่ม แบ่งตามการพรรณนาลักษณะข้อมูล ได้แก่ ค่ากลางของข้อมูล ค่าการแจกแจงของข้อมูล และค่าแสดงตำแหน่งสัมพัทธ์ของข้อมูล ค่าตัวเลขที่ใช้อธิบายลักษณะของประชากรจะเรียกว่า พารามิเตอร์ (parameter) ระบุด้วยตัวอักษรภาษากรีก ส่วนค่าตัวเลขที่ใช้อธิบายลักษณะของตัวอย่าง จะเรียกว่า ค่าสถิติ (statistic) ระบุด้วยตัวอักษรภาษาอังกฤษ

ก. ค่ากลาง คือค่าที่ใช้เป็นตัวแทนสรุปลักษณะแนวโน้มเข้าสู่ส่วนกลางของข้อมูล มี 3 ชนิด

1. ค่าเฉลี่ย (mean: $\bar{x}, \mu$) เป็นค่ากลางที่นิยมใช้มากที่สุด คำนวณได้จากผลรวมของข้อมูล (x_i) ทุกข้อมูลหารด้วยจำนวนข้อมูล (n หรือ N) ดังสมการ 1.4 และ 1.5

$$\text{ค่าเฉลี่ยตัวอย่าง} \quad \bar{x} = \frac{\sum_{i=1}^n x_i}{n} \quad \dots\dots\dots 1.4$$

$$\text{ค่าเฉลี่ยประชากร} \quad \mu = \frac{\sum_{i=1}^N x_i}{N} \quad \dots\dots\dots 1.5$$

2. ค่ามัธยฐาน (median) เป็นค่าที่อยู่ตรงกลางของข้อมูลที่เรียงจากน้อยไปมาก ถ้าจำนวนข้อมูลทั้งหมด (n) เป็นเลขคู่ให้ใช้ค่าเฉลี่ยของค่าที่ $\frac{n}{2}$ และค่าที่ $\frac{n}{2}+1$ การคำนวณค่ามัธยฐานทำได้ง่ายและรวดเร็ว จึงจัดเป็นค่าตัวเลขที่ใช้ใน EDA และสถิติที่ไม่ใช้พารามิเตอร์ (ข้อ 1.3.2)
3. ค่าฐานนิยม (mode) เป็นข้อมูลที่มีความถี่สูงสุด เพราะเป็นค่าที่พบซ้ำมากที่สุด ใช้เป็นตัวแทนของข้อมูลแบบกลุ่ม หรือข้อมูลแบบมีช่วง

การเลือกใช้ค่ากลางขึ้นกับลักษณะการแจกแจงของข้อมูล ถ้าข้อมูลมีการแจกแจงแบบปกติ (normal distribution) หรือใกล้เคียงกับแบบปกติ ให้ใช้ค่ากลางใดก็ได้ในสามค่าเป็นตัวแทนของข้อมูล เนื่องจากการแจกแจงแบบปกติ ค่ากลางทั้งสามชนิดจะมีค่าเท่ากัน แต่ถ้าเกิดการแจกแจงแบบเบ้ซ้ายหรือเบ้ขวา ให้ใช้ค่ามัธยฐาน เนื่องจากถ้าข้อมูลมีการแจกแจงแบบเบ้ซ้าย ค่าเฉลี่ยจะน้อยกว่าค่ามัธยฐาน ในทางตรงกันข้าม ถ้า

ข้อมูลมีการแจกแจงแบบเบ้ขวา ค่าเฉลี่ยจะมากกว่าค่ามัธยฐาน ดังนั้นการเปรียบเทียบค่าเฉลี่ยและค่ามัธยฐานสามารถบอกการแจกแจงของข้อมูลได้ว่าเป็นแบบปกติหรือไม่

ข. ค่าการแจกแจง พารามิเตอร์นี้ใช้ระบุการแจกแจง หรือการกระจายตัวของข้อมูลรอบค่ากลาง

1. พิสัย (range: R) บอกช่วงของการแจกแจงของข้อมูล ว่ามีความกว้างเท่าใด โดยเอาข้อมูลที่มีค่าสูงสุด

($x_{\max}$) ตั้งและลบด้วยข้อมูลที่มีค่าต่ำสุด ($x_{\min}$) ตามสูตร $R = x_{\max} - x_{\min}$

2. ส่วนเบี่ยงเบนมาตรฐาน (standard deviation: s, σ) คือ ค่าเฉลี่ยของส่วนเบี่ยงเบน (ผลต่าง) ระหว่างข้อมูลแต่ละค่า (x_i) กับค่าเฉลี่ยของข้อมูลนั้น ซึ่งผลต่างที่ได้มีทั้งค่าที่เป็นบวกและเป็นลบ จึงยกกำลังสองผลต่างก่อนนำไปหาค่าเฉลี่ย ดังนั้นคำตอบสุดท้ายต้องผ่านการถอดรากที่สอง ออกมาเป็นส่วนเบี่ยงเบนมาตรฐานซึ่งมีเครื่องหมายเป็นได้ทั้งค่าบวกและลบ ดังสมการที่ 1.6 และ 1.7

$$\text{ส่วนเบี่ยงเบนมาตรฐานของตัวอย่าง} \quad s = \pm \sqrt{\frac{\sum (x_i - \bar{x})^2}{n - 1}} \quad \dots\dots\dots 1.6$$

$$\text{ส่วนเบี่ยงเบนมาตรฐานของประชากร} \quad \sigma = \pm \sqrt{\frac{\sum (x_i - \mu)^2}{N}} \quad \dots\dots\dots 1.7$$

เนื่องจากมีเครื่องหมายบวกและลบติดอยู่อย่างนี้ จึงต้องยกกำลังสองให้เครื่องหมายบวกและลบหายไปก่อนนำไปทำการคำนวณ จากนั้นจึงถอดรากที่สองกลับออกมาเป็นส่วนเบี่ยงเบนมาตรฐาน

3. ความแปรปรวน (variance, s^2, σ^2) คือ ค่าส่วนเบี่ยงเบนมาตรฐานยกกำลังสอง ใช้ในการเปรียบเทียบการแจกแจงของข้อมูล 2 ชุด

ค. ค่าแสดงตำแหน่งสัมพัทธ์ ค่านี้นิยามว่าข้อมูลที่สนใจอยู่ตำแหน่งไหนเมื่อเปรียบเทียบกับข้อมูลทั้งหมด

1. เปอร์เซ็นไทล์ (percentile, P) คือการเรียงข้อมูลจากน้อยไปมาก แล้วแบ่งจำนวนข้อมูลออกเป็น 100 ส่วน ถ้า n เป็นจำนวนข้อมูล เปอร์เซ็นไทล์ที่ 50 คือ ข้อมูลที่ตกอยู่ในตำแหน่งที่ $0.5(n+1)$ ความหมายคือ ถ้าเราสอบได้คะแนนเปอร์เซ็นไทล์ ที่ 50 หมายความว่า มีคน 50% ที่สอบได้คะแนนน้อยกว่าเรา เปอร์เซ็นไทล์ที่ 50 ก็คือค่ามัธยฐาน สูตรสำหรับหาเปอร์เซ็นไทล์ คือ $P = (P/100)(n+1)$

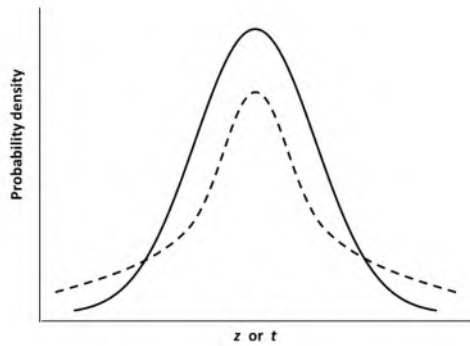
2. ควอไทล์ (quartile, Q) คือการเรียงข้อมูลจากน้อยไปมาก แล้วแบ่งจำนวนข้อมูลออกเป็น 4 ส่วน ข้อมูลตกอยู่ในส่วนที่เท่าไร ก็เรียกว่าควอไทล์ที่เท่าไรนั้น ถ้า n เป็นจำนวนข้อมูล ควอไทล์ที่ 1 คือ ข้อมูลที่ตกอยู่ในตำแหน่งที่ $(n+1)/4$ ดังนั้นค่าควอไทล์ที่ 1 (Q_1) ก็คือค่าเปอร์เซ็นไทล์ที่ 25 ช่วงกว้างระหว่าง $Q_3 - Q_1$ เรียกว่า IQR (Interquartile range) หรือ ค่าพิสัยควอไทล์

3. คะแนนมาตรฐาน Z (Z -standard score) เป็นค่าความแตกต่างระหว่างข้อมูลและค่าเฉลี่ย (μ) ของประชากร โดยคิดเป็นจำนวนเท่าของส่วนเบี่ยงเบนมาตรฐาน (σ) ของประชากร ตามสมการ

$$Z = \frac{x - \mu}{\sigma} \text{ โดย } Z \text{ เป็นแกน } x \text{ ของกราฟการแจกแจงความถี่แบบเส้นโค้งปกติมาตรฐาน ที่มีค่า } \mu =$$

0 และ $\sigma = 1$ ดังในรูปที่ 1.5 พื้นที่ใต้เส้นโค้งคือค่าโอกาสที่จะพบข้อมูลที่ตำแหน่ง Z บนแกน x

สามารถเปิดดูได้จากตารางค่า Z ซึ่งพื้นที่ใต้เส้นโค้งปกติมาตรฐานนี้จะมีค่าเท่ากับพื้นที่ที่คำนวณจากเส้นโค้งปกติทั่วไปที่มีค่า μ และ σ ต่างกันตามที่เกิดขึ้นจริงของประชากรใดประชากรหนึ่ง ตารางค่า Z นี้จึงมีประโยชน์ในการใช้เปิดตารางหาค่าพื้นที่ใต้โค้งหรือค่าความน่าจะเป็นที่จะเกิดขึ้นของข้อมูล



นั้นในประชากรที่ศึกษา กรณีที่ไม่ทราบค่า σ แต่ทราบค่าจริง ซึ่งเทียบเท่าเป็นค่า μ จะใช้คะแนนมาตรฐาน t (t -standard score) ตามสมการ $t = \frac{x - \mu}{s/\sqrt{n}}$ โดย t เป็นแกน x ของกราฟการแจกแจงความถี่แบบที่ ดังในรูปที่ 1.5

รูปที่ 1.5 กราฟการแจกแจงปกติ (เส้นทึบ) และการแจกแจงแบบที (เส้นประ)

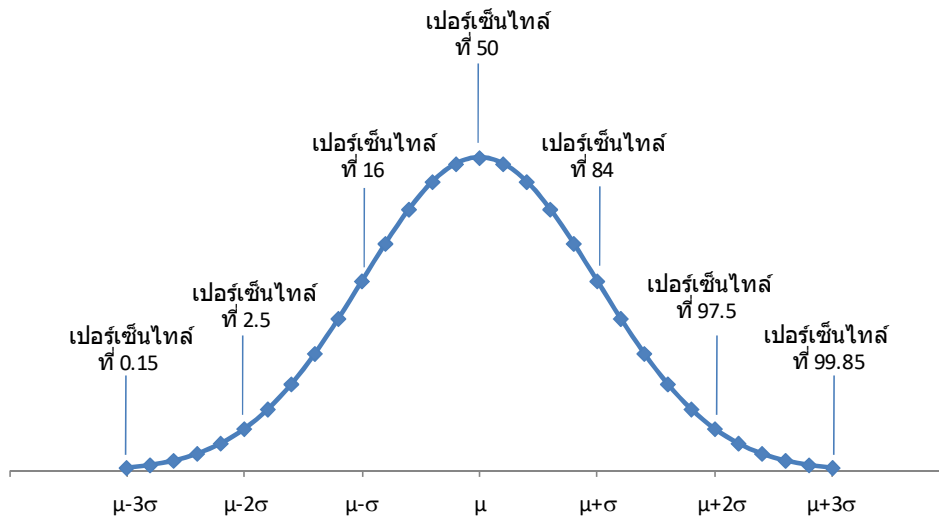
1.2 การแจกแจงแบบปกติ (normal distribution)

เป็นการแจกแจงความน่าจะเป็นแบบต่อเนื่องที่มีสมมาตร ข้อมูลจะมีแนวโน้มเกาะกลุ่มหนาแน่นอยู่ตรงกลาง ซึ่งมีค่าเฉลี่ยอยู่ เป็นข้อมูลที่พบได้บ่อยที่สุด มีความถี่สูงสุด และลดลงสองข้างเป็นจำนวนเท่าของส่วนเบี่ยงเบนมาตรฐาน (σ) ไปทางด้านที่ข้อมูลมีค่ามากกว่าและน้อยกว่าค่าเฉลี่ย (μ) โดย 68% ของประชากรจะมีค่าอยู่ระหว่างช่วง $\pm 1\sigma$ ของค่าเฉลี่ย 95% ของประชากรจะมีค่าอยู่ระหว่างช่วง $\pm 2\sigma$ ของค่าเฉลี่ย 99.7% ของประชากรจะมีค่าอยู่ระหว่างช่วง $\pm 3\sigma$ ของค่าเฉลี่ย ดังนั้นเมื่อเปลี่ยนค่าเปอร์เซ็นต์เหล่านี้ให้เป็นค่าเปอร์เซ็นต์ไทล์ จะได้รายละเอียดดังในตารางที่ 1.1 (Glantz, 2012) และนำไปพลอตกราฟได้ดังในรูปที่ 1.6

ตารางที่ 1.1 การคำนวณตำแหน่งบนเส้นโค้งปกติในรูปเปอร์เซ็นต์ไทล์

ตำแหน่งบน เส้นโค้งปกติ	จำนวนประชากร นับจากค่ากลาง	จำนวนประชากร นับจากด้านซ้าย	ตำแหน่งเปอร์เซ็นต์ไทล์ บนเส้นโค้งปกติ (P)
$\mu - 3\sigma$	99.7/2 = 49.85%	50-49.85 = 0.15%	P 0.15
$\mu - 2\sigma$	95/2 = 47.50%	50-47.50 = 2.50%	P 2.5
$\mu - \sigma$	68/2 = 34%	50-34 = 16%	P 16
μ	50%	50%	P 50
$\mu + \sigma$	68/2 = 34%	50+34 = 84%	P 84
$\mu + 2\sigma$	95/2 = 47.50%	50+47.5 = 97.50%	P 97.5
$\mu + 3\sigma$	99.7/2 = 49.85%	50+49.85 = 99.85%	P 99.85

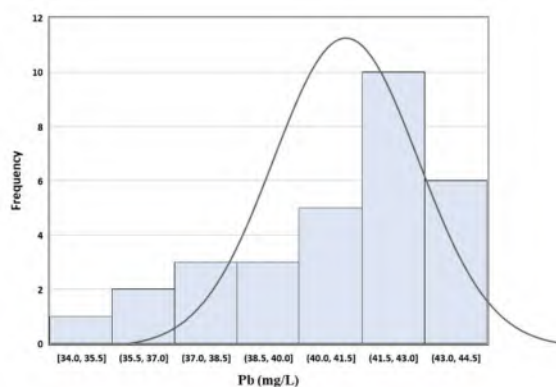
ดังนั้นในการทดสอบว่าข้อมูลมีการแจกแจงแบบปกติหรือไม่ ให้เรียงข้อมูลจากน้อยไปหามาก และหาค่าเปอร์เซ็นต์ไทล์ที่ P 0.15 ถึง P 99.85 ถ้ามีค่าเท่ากับ $\mu - 3\sigma$ ถึง $\mu + 3\sigma$ ตามลำดับ แสดงว่าการแจกแจงของข้อมูลชุดนี้เป็นแบบปกติ



รูปที่ 1.6 ตำแหน่งเปอร์เซ็นต์ไทล์บนเส้นโค้งปกติ

นอกจากนี้ยังมีวิธีตรวจสอบว่าข้อมูลมีการแจกแจงแบบปกติอีก 2 วิธี ที่เป็นการตรวจสอบอย่างง่ายและรวดเร็วดังนี้ (Charles Zaiontz, [www. https://real-statistics.com/](https://real-statistics.com/))

1.2.1 วิธีกราฟิก นอกจากวิธีพลอตกราฟ cumulative distribution และ normality probability plot ที่ได้กล่าวไปแล้ว ยังมีการสร้างฮิสโทแกรม และ box plots ซึ่งจะทำให้สังเกตเห็นแนวโน้มของข้อมูลว่ามีลักษณะการแจกแจงปกติหรือไม่ โดยการนำเส้นโค้งปกติในรูปที่ 1.5 ไปวางทับบนฮิสโทแกรมของข้อมูล จะพบว่าข้อมูลเบ้าซ้าย ไม่มีการแจกแจงแบบปกติดังในรูปที่ 1.7



รูปที่ 1.7 เส้นโค้งปกติวางทับบนฮิสโทแกรม

ส่วน box plot ไม่สามารถดูลักษณะโค้งปกติได้ แต่ใช้ดูความสมมาตรของข้อมูลซึ่งบางครั้งก็เพียงพอที่จะใช้แทนการแจกแจงแบบปกติในการพิจารณาการใช้สถิติทดสอบบางชนิดได้

1.2.2 การวิเคราะห์ค่าความเบ้ (skewness) และความโด่ง (kurtosis)

ถ้าค่าความเบ้ และความโด่ง ต่างก็ไม่เกิน 2 เท่าของ ค่าความคลาดเคลื่อนมาตรฐานของความเบ้ และความโด่ง จะถือว่าข้อมูลชุดนั้นมีการแจกแจงแบบปกติ

โดยค่าความคลาดเคลื่อนมาตรฐานของความเบ้มีค่าประมาณ $\sqrt{6/n}$ และ n คือ ขนาดของตัวอย่าง และค่าความคลาดเคลื่อนมาตรฐานของความโด่งมีค่าประมาณ $\sqrt{24/n}$ และ n คือ ขนาดของตัวอย่าง โดยค่าความเบ้และความโด่งของข้อมูลในรูปที่ 1.1 ซึ่งแสดงดังในตารางข้างล่างนี้

	A	B	C	D	E	F	G	H	I	J
1	39.00	39.10	39.20	39.40	39.50	39.60	39.70	39.80	39.80	39.90
2	39.90	39.90	39.90	40.00	40.00	40.00	40.00	40.10	40.10	40.20
3	40.20	40.30	40.30	40.40	40.50	40.60	40.60	40.70	40.80	41.00

สามารถคำนวณได้โดยการใช้ฟังก์ชันในโปรแกรมเอ็กเซล ดังนี้

ความเบ้ = SKEW(A1:J3) = -0.155 ความโด่ง = KURT(A1:J3) = -0.148 และ $n = 30$

ดังนั้น ความคลาดเคลื่อนมาตรฐานของความเบ้ = $\sqrt{6/30} = 0.447$

ค่าความคลาดเคลื่อนมาตรฐานของความโด่ง = $\sqrt{24/30} = 0.894$

นำค่าความเบ้ที่ได้ไปเปรียบเทียบกับความคลาดเคลื่อนมาตรฐานโดยไม่ดูเครื่องหมาย จะได้ว่า ความเบ้ $0.155 < 2 \times 0.447$ และ ความโด่ง $0.148 < 2 \times 0.894$ ข้อมูลชุดนี้จึงมีการแจกแจงแบบปกติ เพราะมีความเบ้และความโด่งน้อยกว่า 2 เท่าของค่าความคลาดเคลื่อนมาตรฐานของความเบ้และความโด่ง

นอกจากสองวิธีนี้แล้วยังมีการทดสอบด้วยวิธีทางสถิติ เช่น Chi-square Test, Kolmogorov-Smirnov (KS) Test, Lilliefors Test, Shapiro-Wilk (SW) Test, Jarque-Barre Test และ D'Agostino-Pearson Test ซึ่งสามารถศึกษาเพิ่มเติมได้ที่ [www. https://real-statistics.com/](https://real-statistics.com/)

1.3 สถิติเชิงอนุมาน (inferential or inductive statistics)

คือ สถิติที่ใช้ในการวิเคราะห์ข้อมูลที่ได้จากกลุ่มตัวอย่าง เพื่อบอกลักษณะของประชากรที่สุ่มตัวอย่างมา หรือเพื่อเปรียบเทียบข้อมูลจากประชากรสองกลุ่ม โดยมีการนำทฤษฎีความน่าจะเป็นมาใช้ในการวิเคราะห์ข้อมูล สถิติเชิงอนุมานนี้ แบ่งออกเป็น 2 ประเภท คือ สถิติที่ใช้พารามิเตอร์ และสถิติที่ไม่ใช้พารามิเตอร์

1.3.1 สถิติที่ใช้พารามิเตอร์ (parametric statistics)

เป็นสถิติที่มีพารามิเตอร์มาเกี่ยวข้องในการคำนวณเพื่อประมาณค่า ทดสอบ หรือวิเคราะห์ ใช้กับข้อมูลเชิงปริมาณและประชากรมีการแจกแจงแบบปกติ วิธีทางสถิติประเภทนี้ได้แก่ การประมาณค่าพารามิเตอร์ การทดสอบสมมติฐานทางสถิติ และการวิเคราะห์การถดถอยและสหสัมพันธ์ ซึ่งแตกต่างกันดังนี้

- การประมาณค่าพารามิเตอร์

การประมาณค่าพารามิเตอร์มี 2 แบบ คือ การประมาณแบบค่าเดียว (point estimation) และการประมาณแบบเป็นช่วง (interval estimation) การประมาณค่าพารามิเตอร์ ที่สำคัญ ได้แก่ การประมาณค่าเฉลี่ยประชากร (μ) ด้วยค่าเฉลี่ยตัวอย่าง ($\bar{X}$) การประมาณค่าส่วนเบี่ยงเบนมาตรฐานประชากร (σ) ด้วยค่าส่วนเบี่ยงเบนมาตรฐานตัวอย่าง (s)

- การทดสอบสมมติฐาน (hypothesis testing)

เมื่อมีความสงสัยในเรื่องใด เราจะตั้งสมมติฐานเกี่ยวกับปัญหานั้น ซึ่งก็คือการตั้งสมมติฐานเกี่ยวกับประชากร จากนั้นก็ทำการเก็บข้อมูลจากกลุ่มตัวอย่าง เพื่อหาข้อสรุปว่า สมมติฐานที่ตั้งนั้น เราจะยอมรับหรือปฏิเสธ ทางเคมีวิเคราะห์จะเรียกว่า การทดสอบนัยสำคัญ (significant test) เนื่องจากความสงสัยที่เกิดขึ้นจะเกี่ยวกับผลการวิเคราะห์ ว่าต่างกันอย่างมีนัยสำคัญหรือไม่ หรือแตกต่างเพราะความคลาดเคลื่อนแบบสุ่มที่ไม่สามารถตัดทิ้งได้

- การวิเคราะห์การถดถอยและสหสัมพันธ์ (regression and correlation analysis)

การถดถอยเป็นการหาความสัมพันธ์ของตัวแปรตาม (แกน Y) และตัวแปรต้น (แกน X) ดูแนวโน้มว่าเมื่อเปลี่ยนตัวแปรต้น แล้วตัวแปรตามจะเพิ่มหรือลด เปลี่ยนไปเท่าไร ค่าสัมประสิทธิ์สหสัมพันธ์ (correlation coefficient) มีค่าตั้งแต่ 0-1 เป็นค่าที่บอกเชิงปริมาณของความสัมพันธ์ที่ทำได้ ว่าเป็นไปตามความสัมพันธ์นั้นมากหรือน้อยแค่ไหน โดยแบ่งออกได้เป็น 4 ระดับจากน้อยไปมาก ดังนี้

0.00-0.25	ไม่มีความสัมพันธ์ ถึงมีความสัมพันธ์น้อยมาก
0.25-0.50	มีความสัมพันธ์ น้อยมาก ถึงปานกลาง
0.50-0.75	มีความสัมพันธ์ ปานกลางถึงดี
0.75-1.00	มีความสัมพันธ์ ดีถึงดีมาก

1.3.2 สถิติที่ไม่ใช้พารามิเตอร์ (non-parametric statistics)

สถิติที่ไม่ใช้พารามิเตอร์ เป็นสถิติที่ใช้กับข้อมูลที่ไม่มีความต่อเนื่องเกี่ยวกับการแจกแจงของข้อมูล ไม่ว่าจะเป็นแบบไหน อย่างไร ก็ใช้วิธีนี้ได้ จึงมีชื่อเรียกอีกอย่างว่า สถิติที่มีการแจกแจงอิสระ (distribution free statistics) เป็นสถิติที่ไม่ได้ทดสอบเกี่ยวกับค่าพารามิเตอร์ของประชากร สามารถใช้ได้กับข้อมูลที่มีการวัดระดับต่ำสุด คือ ข้อมูลลักษณะกลุ่มหรือประเภท (nominal or category data) ข้อดีของวิธีนี้คือ ใช้การคำนวณอย่างง่าย สามารถคิดในใจได้ถ้าข้อมูลมีจำนวนน้อยมาก โดยมีวิธีทางสถิติเทียบได้กับวิธีทางสถิติที่ใช้พารามิเตอร์ ได้แก่ การประมาณค่าทางสถิติ การทดสอบสมมติฐานทางสถิติ การวิเคราะห์การถดถอยและสหสัมพันธ์ เป็นต้น โดยภาพรวมจะไม่ค่อยพบว่ามีการใช้สถิติที่ไม่ใช้พารามิเตอร์ ในทางเคมีวิเคราะห์มากเท่าไร อาจเป็นเพราะสามารถเพิ่มจำนวน n ได้โดยการซ้ำในห้วงปฏิบัติการทำให้ได้ข้อมูลมีการแจกแจงแบบปกติ หรือเป็นเพราะโดยธรรมชาติข้อมูลที่ได้เป็นการตรวจวัดปริมาณของสารในตัวอย่างไม่มีลักษณะที่เป็น การแจกแจงแบบปกติอยู่แล้ว (Miller & Miller, 2010) ดังนั้นในที่นี้จะกล่าวแทรกถึงสถิติที่ไม่ใช้พารามิเตอร์เพียงบางส่วน ไว้ในบางบทที่มีการใช้ ได้แก่ บทที่ 6 การใช้สถิติ Wilcoxon sign rank test ในการเปรียบเทียบ

ข้อมูล 2 ชุดว่าต่างกันหรือไม่เพื่อตรวจสอบความถูกต้อง (accuracy) ของวิธี และบทที่ 7 สถิติเชิงปรับแก้ (robust analysis) ที่ใช้กับข้อมูลที่มีค่าสุดต่าง ในการทดสอบความชำนาญ

1.4 การใช้สถิติในงานทางเคมีวิเคราะห์

เคมีวิเคราะห์เป็นสาขาหนึ่งของวิชาเคมี ที่เน้นศึกษาเกี่ยวกับการวิเคราะห์หาปริมาณของสาร เช่น การหาปริมาณตะกั่วในอากาศที่เก็บจากถนนที่มีการจราจรคับคั่ง การหาปริมาณน้ำตาลในเลือดของผู้ป่วย การหาปริมาณไซยาไนด์ในน้ำทิ้งจากโรงงานอุตสาหกรรมซุโหละ การทดสอบคุณภาพน้ำดื่ม การทดสอบโลหะผสม เพื่อยืนยันคุณภาพก่อนนำไปใช้ในอุตสาหกรรม การตรวจสอบทางนิติวิทยาศาสตร์เพื่อใช้ในการดำเนินคดีอาชญากรรม เป็นต้น วิธีวิเคราะห์หาปริมาณนี้แบ่งได้ออกเป็น 2 วิธีคือ การวิเคราะห์ปริมาณแบบดั้งเดิม เป็นการไทเทรต (titrimetric method) และการวิเคราะห์เชิงน้ำหนัก (gravimetric method) ส่วนอีกวิธีซึ่งเป็นที่นิยมมากในปัจจุบันคือการวิเคราะห์โดยใช้เครื่องมือทางวิทยาศาสตร์ ไม่ว่าจะวิเคราะห์ปริมาณด้วยวิธีใด โดยหลักการวัดค่าทุกชนิดทางวิทยาศาสตร์ เช่น น้ำหนัก ปริมาตร อุณหภูมิ ความดัน ความยาว ค่าที่วัดได้จะมีความไม่แน่นอนปรากฏอยู่ด้วยที่ตำแหน่งสุดท้ายของค่าที่วัดได้เสมอ (uncertainty in the last digit of measurement) ไม่ว่าจะใช้เครื่องมือวัดค่าที่ละเอียดแค่ไหน การวิเคราะห์ปริมาณจึงมักทำซ้ำหลายครั้งสำหรับสารตัวอย่างชนิดเดียวกัน (repeated measurement) เพื่อให้เกิดความน่าเชื่อถือของผลที่ได้ โดยคำตอบที่ได้จะไม่ตอบเป็นค่าใดค่าหนึ่งค่าเดียว แต่จะต้องมีการประมาณค่าความคลาดเคลื่อนมาด้วยเสมอ คำตอบที่ได้ในการวิเคราะห์ปริมาณจึงต้องอาศัยวิธีทางสถิติคำนวณตอบออกมาเป็นค่าทางสถิติ เช่น ใช้ค่าเฉลี่ย (mean) เป็นคำตอบ และมีค่าส่วนเบี่ยงเบนมาตรฐาน (standard deviation) มาคำนวณเป็นค่าแสดงความคลาดเคลื่อน หรือเป็นค่าแสดงความเที่ยง (precision) ซึ่งบอกถึงการแจกแจงของคำตอบที่ได้ เป็นต้น

นอกจากนี้การวิเคราะห์ปริมาณโดยใช้เครื่องมือทางวิทยาศาสตร์ มีการสร้างกราฟมาตรฐาน (calibration curve) ซึ่งส่วนใหญ่จะเป็นกราฟเส้นตรง โดยการวัดสัญญาณของสารมาตรฐานที่รู้ความเข้มข้นด้วยเครื่องมือที่เลือกใช้ นำค่าที่ได้มาพลอตกราฟ มีแกน y เป็นสัญญาณ และแกน x เป็นความเข้มข้น สมการแสดงความสัมพันธ์ของสัญญาณและความเข้มข้นหาได้โดยใช้วิธีทางสถิติที่เรียกว่าการวิเคราะห์ความถดถอย (regression analysis) การพัฒนาวิธีวิเคราะห์โดยเครื่องมือ หรือพัฒนาเครื่องมือขึ้นมาใหม่ จะต้องมีการทดสอบความถูกต้องของผลที่ได้ โดยการเปรียบเทียบซึ่งใช้วิธีทางสถิติเช่นกัน เช่น ใช้ t -test ในการเปรียบเทียบค่าเฉลี่ยของข้อมูล 2 ชุด ใช้ F -Test ใช้ในการเปรียบเทียบความเที่ยง ของข้อมูล 2 ชุด โดยภาพรวมสามารถสรุปการใช้สถิติทางเคมีวิเคราะห์ออกเป็น 2 ประเภท ได้แก่ การใช้สถิติในการรายงานผล (รวมทั้งการสร้างกราฟมาตรฐาน) และการใช้สถิติในการเปรียบเทียบผล นอกจากนี้ในการควบคุมคุณภาพการวิเคราะห์ในห้องปฏิบัติการก็มีการใช้สถิติเช่นกัน สำหรับการวิเคราะห์ในการสุ่มตัวอย่างทางเคมีวิเคราะห์ จะใช้มากโดยเฉพาะกรณีของการนำผลวิเคราะห์ไปใช้ทางสิ่งแวดล้อม

1.5 ความคลาดเคลื่อนในการทำปริมาณวิเคราะห์

ความคลาดเคลื่อนที่เกิดขึ้นในการวิเคราะห์ปริมาณสามารถแบ่งออก เป็น 2 ชนิด คือ ความคลาดเคลื่อนแบบมีระบบ (systematic error) และความคลาดเคลื่อนแบบสุ่ม (random error) นอกจากนี้หากผู้ทำการทดลองมีความประมาท ทำให้เกิดความผิดพลาด ซึ่งเกิดเป็นบางครั้ง เช่น การเตรียมสารมาตรฐานผิดพลาด ความเข้มข้นมากหรือน้อยกว่าค่าจริง ทำสารตัวอย่างหกหล่น สารรีเอเจนต์ที่ใช้มีการปนเปื้อน หรือเครื่องมือเกิดเสีระหว่างการทำวิเคราะห์เนื่องจากการไม่มีการซ่อมบำรุงที่ดี ความคลาดเคลื่อนที่เกิดขึ้นนี้จะมีค่ามาก ผลการวิเคราะห์จะแตกต่างอย่างค่อนข้างชัด ความคลาดเคลื่อนชนิดนี้ เรียกว่า ความคลาดเคลื่อนแบบรวบยอด (gross error) หากเกิดความคลาดเคลื่อนแบบนี้ขึ้น จะไม่มีทางเลือกอื่นใด นอกจากทิ้งผลการวิเคราะห์นั้นเสีย และ เริ่มต้นทำการวิเคราะห์ใหม่ตั้งแต่แรก ความคลาดเคลื่อนแบบรวบยอดเป็นความคลาดเคลื่อนที่สังเกตและตรวจสอบได้ง่าย และที่สำคัญหลีกเลี่ยงได้ถ้ามีความระมัดระวัง ต่างจากความคลาดเคลื่อน 2 ชนิดที่กล่าวมา ซึ่งสังเกตและแยกประเภทได้ยากกว่า แต่ถ้าไม่แน่ใจว่าผลที่ค่อนข้างแตกต่างไปจากกลุ่มมาจากความคลาดเคลื่อนแบบรวบยอดหรือไม่ ก็มีการวิเคราะห์ทางสถิติที่เรียกว่า การทดสอบ outlier ไว้สำหรับตรวจสอบเพื่อจะตัดค่านี้ออกหรือไม่

- **ความคลาดเคลื่อนแบบมีระบบ** (systematic error) เป็นความคลาดเคลื่อนที่ควบคุมได้ แก้ไขได้ มีค่าแน่นอน มีแนวโน้มไปในทางเดียวกัน เป็นได้ทั้งทางบวกหรือทางลบ แต่จะไม่เกิดเป็นแบบทางบวกหรือทางลบสลับกันไป

ในการวิเคราะห์แต่ละครั้งอาจมีความคลาดเคลื่อนแบบมีระบบจากหลายสาเหตุต่างกัน บางสาเหตุไปทางบวก บางสาเหตุไปทางลบ ผลรวมค่าความคลาดเคลื่อนแบบมีระบบทั้งหมดจากทุกสาเหตุ เรียกว่า ค่าเอนเอียง (bias)

ตัวอย่างสาเหตุของความคลาดเคลื่อนแบบมีระบบ ได้แก่ การใช้เครื่องชั่งหรือเครื่องมือที่ไม่ได้สอบเทียบ (calibration) อย่างถูกวิธีหรือ สอบเทียบไม่ตรงตามกำหนดเวลา (ไม่สม่ำเสมอ ทั้งช่วงนานไป) การใช้สารเคมีที่ความบริสุทธิ์ไม่ตรงตามที่ระบุ ปฏิบัติการเคมีที่เกิดซ้ำหรือเกิดไม่สมบูรณ์ สารเคมีบางตัวที่ไม่เสถียร ความคลาดเคลื่อนแบบมีระบบนี้ทำให้ผลการวิเคราะห์ที่ได้แตกต่างไปจากค่าจริงที่ควรจะเป็น เรียกว่ามีผลต่อค่าความถูกต้อง (accuracy) ของผลการวิเคราะห์ การตรวจสอบความถูกต้อง ใช้สถิติในการเปรียบเทียบผล เช่น ใช้ t -test เปรียบเทียบผลที่ได้จากการวิเคราะห์ด้วยวิธีที่พัฒนาขึ้นมาใหม่ กับผลที่วิเคราะห์ด้วยวิธีมาตรฐาน หรือใช้วิเคราะห์ด้วยวิธีที่พัฒนาขึ้นมาใหม่นี้ไปตรวจสอบว่าสอดคล้องรับรองและเปรียบเทียบผลกับค่าจริงที่เป็นค่ารับรองมาพร้อมกับวัสดุอ้างอิงนั้น ความคลาดเคลื่อนแบบมีระบบนี้ บางครั้งเรียกว่าความคลาดเคลื่อนที่มาจากปัจจัยภายใน ซึ่งหมายถึงปัจจัยที่เกี่ยวข้องกับวิธีวิเคราะห์ และสามารถแก้ไขได้

- **ความคลาดเคลื่อนแบบสุ่ม** (random error) เป็นความคลาดเคลื่อนที่ควบคุมไม่ได้ แก้ไขไม่ได้ ไม่มีทิศทางแน่นอน จะเกิดแบบไปทางบวกหรือทางลบสลับกันไปมา เช่น การอ่านค่าที่อยู่ระหว่างสเกลด้วยสายตา

ของแต่ละคน เป็นค่าที่กะประมาณเอา การปล่อยสารให้ไหลออกจากปิเปตจนหมด ใช้เวลาต่างกัน หรือองศาของปิเปตที่เอียงต่างกันในการปล่อยสาร ความแปรปรวนของอุณหภูมิทำให้ความหนืดและปริมาตรของของเหลวเปลี่ยน สัญญาณรบกวนของเครื่องมือบางชนิดที่มาจากระบบไฟฟ้าซึ่งเรียกว่า background noise ซึ่งได้จากการวัดสารละลายแบลนด์ การทำงานของเครื่องชั่ง ค่าตำแหน่งสุดท้ายของน้ำหนักที่ได้จากการชั่ง จะมีความไม่แน่นอนเสมอ มาจากการสั่นสะเทือนหรือ ลมพัดซึ่งเป็นสาเหตุที่ไม่สามารถแก้ไขให้หายไปได้ ค่าความไม่แน่นอนที่ปรากฏอยู่ที่ตำแหน่งสุดท้ายของค่าที่วัดได้ด้วยเครื่องมือวัด เป็นค่าที่กะประมาณเอาตามหลักของการคิดเลขนัยสำคัญ ค่านี้อาจเป็นสาเหตุหนึ่งของความคลาดเคลื่อนแบบสุ่มที่อาจกล่าวได้ว่ามีอยู่ในเกือบทุกการวิเคราะห์ปริมาณที่มีการใช้อุปกรณ์และเครื่องมือ นอกจากนี้ยังมีปัจจัยภายนอกที่ไม่เกี่ยวกับวิธีวิเคราะห์ เช่น การเปลี่ยนผู้วิเคราะห์ การเปลี่ยนวันที่วิเคราะห์ การเปลี่ยนห้องปฏิบัติการวิเคราะห์ ความแปรปรวนของอุณหภูมิและความดันในห้องปฏิบัติการ การใช้สารเคมีจากผู้ผลิตคนละแห่งหรือคนละล็อตของการผลิต เป็นต้น ความคลาดเคลื่อนแบบสุ่มนี้มีผลต่อค่า ความเที่ยง (precision) ของผลการวิเคราะห์ ซึ่งสามารถรายงานด้วยค่าส่วนเบี่ยงเบนมาตรฐานสัมพัทธ์ (relative standard deviation)

ความคลาดเคลื่อนทั้งสองชนิดนี้ สามารถเปรียบเทียบกันได้ตามตารางที่ 1.2

ตารางที่ 1.2 ลักษณะของความคลาดเคลื่อนแบบมีระบบและแบบสุ่ม (Miller & Miller, 2010)

ความคลาดเคลื่อนแบบมีระบบ	ความคลาดเคลื่อนแบบสุ่ม
ทำให้เกิดความเอนเอียง (bias) มีผลต่อค่าความถูกต้อง (accuracy)	มีผลต่อค่าความเที่ยง (precision) หรือ repeatability และ reproducibility
ควบคุมได้ แก้ไขได้ โดยเปรียบเทียบกับวิธีมาตรฐาน หรือวัสดุอ้างอิงรับรอง	ควบคุมไม่ได้ แก้ไขไม่ได้ ไม่สามารถทำให้เป็นศูนย์ได้ แต่ทำให้ลดลงได้ด้วยการวิเคราะห์อย่างระมัดระวัง ด้วยเทคนิคที่ดีจากผู้ทำการทดลองที่มีประสบการณ์
มีแนวโน้มไปในทางเดียวกัน อาจจะมากกว่าค่าจริง หรือน้อยกว่าค่าจริง	ไม่มีทิศทางแน่นอน ทำให้ผลที่วิเคราะห์ซ้ำ (replicate) แกว่งไปมารอบค่าเฉลี่ย
ไม่สามารถหาค่าได้จากการวิเคราะห์ซ้ำ	หาค่าได้จากการวิเคราะห์ซ้ำ
รายงานเชิงคุณภาพด้วยการทดสอบสมมติฐานว่าต่างหรือไม่ต่างจากค่าจริง	รายงานเชิงปริมาณด้วยค่าส่วนเบี่ยงเบนมาตรฐานสัมพัทธ์
เกิดจากผู้ทดลองหรือจากเครื่องมือ	เกิดจากผู้ทดลอง จากเครื่องมือ หรือจากสิ่งแวดล้อม

แบบฝึกหัดท้ายบทที่ 1

- 1.1 สถิติแบ่งออกเป็นกี่ประเภท อะไรบ้าง และในการทำปริมาณวิเคราะห์ มีการใช้สถิติประเภทใดในขั้นตอนไหนของการทำปริมาณวิเคราะห์ จงอธิบาย
- 1.2 การวิเคราะห์หาปริมาณกรดอะซิติกในน้ำส้มสายชู ที่เก็บตัวอย่างมาจากตลาด จำนวน 100 ตัวอย่าง ได้ผลวิเคราะห์ดังนี้ ให้สร้างฮิสโทแกรมของข้อมูลชุดนี้ โดยใช้ค่า bin width ที่คำนวณได้จากสมการที่ 1.1 1.2 และ 1.3 เปรียบเทียบฮิสโทแกรมที่ได้ สมการไหนให้จำนวนชั้นใกล้เคียงค่า $\sqrt{n}$

3.7	4.2	4.7	4.0	4.9	4.3	4.0	4.2	4.7	4.1
4.2	3.8	4.2	4.2	3.8	4.7	3.9	4.7	4.1	4.9
4.4	4.2	4.0	4.3	4.1	4.5	3.8	4.8	4.3	4.4
4.4	4.8	4.5	4.3	4.0	4.3	4.5	4.1	4.4	4.2
4.3	4.6	4.5	4.4	3.9	4.6	3.8	4.0	4.3	4.4
4.2	3.9	3.6	4.1	3.8	4.4	4.1	4.2	4.7	4.3
4.4	4.3	4.1	4.0	4.7	4.6	4.3	4.1	4.2	4.6
4.8	4.5	4.3	4.0	3.9	4.4	4.2	4.0	4.1	4.5
4.9	4.8	3.9	3.8	4.0	4.9	4.5	4.7	4.4	4.6
4.4	3.9	4.2	4.6	4.2	4.4	4.4	4.1	4.8	4.0
- 1.3 ให้หาว่าข้อมูลในข้อ 1.2 มีการแจกแจงแบบปกติหรือไม่ โดยใช้การวิเคราะห์ค่าความเบ้ (skewness) และความโด่ง (kurtosis) ในการตรวจสอบ
- 1.4 ค่าทศนิยมตำแหน่งสุดท้ายของน้ำหนักที่ชั่งจากเครื่องชั่งละเอียด มักจะไม่ค่อยนิ่ง เป็นลักษณะข้อมูลที่เกิดจากความคลาดเคลื่อนแบบใด ค่าสถิติที่ใช้บอกลักษณะของข้อมูลนี้ในเชิงปริมาณคือค่าใด
- 1.5 ความคลาดเคลื่อนในการทำปริมาณวิเคราะห์ เกี่ยวข้องกับสถิติอย่างไร

เอกสารอ้างอิง

Miller J.N. & Miller J.C. (2010) *Statistics and chemometrics for analytical chemistry*, 6th ed., Pearson Education Limited, UK.

Thompson, M., & Lowthian, P.J. (2011). *Notes on statistics and data quality for analytical chemists*, 1st ed., Imperial College Press. London, UK.

Glantz, Stanton A. (2012). *Primer of biostatistics*. 7th ed., McGraw-Hill; New York, USA.

ข้อมูลและการตรวจวัดซ้ำ

การหาปริมาณทางเคมีวิเคราะห์ ข้อมูลดิบคือข้อมูลเบื้องต้นที่ได้จากห้องปฏิบัติการ จะเป็นในรูปแบบน้ำหนัก ปริมาตร เวลา อุณหภูมิ และค่าสัญญาณที่ได้จากเครื่องมือวัดทางวิทยาศาสตร์ เช่น ค่าการดูดกลืนแสง (absorbance) ค่าการคายแสง (emission) ค่าความเข้มแสง (intensity) เป็นต้น ซึ่งอาจรายงานในรูปแบบพื้นที่พีค (peak area) ความสูงพีค (peak height) หรือการอ่านค่าในขณะใดขณะหนึ่ง (transient signal) ของสัญญาณที่มาแบบต่อเนื่อง (ไม่มีลักษณะเป็นพีค) จากนั้นข้อมูลดิบเหล่านี้ จะผ่านการคำนวณออกมาเป็นหน่วยความเข้มข้นเพื่อบอกปริมาณ เช่น โมล/ลิตร (mol/L) พีพีเอ็ม (ppm) หรือ มิลลิกรัมต่อกิโลกรัม (mg/kg) พีพีบี (ppb) หรือ ไมโครกรัมต่อกิโลกรัม ($\mu\text{g/kg}$) ข้อมูลขั้นสุดท้ายนี้จะถูกนำมาคำนวณด้วยวิธีทางสถิติเพื่อรายงานผลหรือเปรียบเทียบผล ทางสถิติเรียกข้อมูลเหล่านี้ว่า ตัวแปร (variables) ซึ่งหมายถึงลักษณะหรือสมบัติของประชากรที่ผู้ศึกษาสนใจ

ปัจจุบันเครื่องมือวัดทางวิทยาศาสตร์บางประเภท อ่านวัดความสะอาดให้ผู้วิเคราะห์โดยใส่สมการสำหรับคำนวณไว้ในเครื่อง และให้คำตอบเป็นข้อมูลขั้นสุดท้ายเลย ซึ่งอาจเป็นในรูปแบบน้ำหนักหรือความเข้มข้น แต่มักจะเป็นการวิเคราะห์แบบกึ่งปริมาณวิเคราะห์ (semi-quantitative analysis) คือวัดในระดับปริมาณที่มาก บอกค่าเป็นช่วงไม่ละเอียดนัก เช่น บอกเป็นเปอร์เซ็นต์น้ำหนักต่อปริมาตร (%w/v) เปอร์เซ็นต์น้ำหนักต่อน้ำหนัก (%w/w) เป็นต้น ข้อมูลแบบนี้จึงต่างไปจากข้อมูลดิบที่ได้มาจากการวิเคราะห์แบบดั้งเดิม ด้วยการไทเทรต หรือการตกตะกอน ข้อมูลดิบเหล่านี้ต้องนำมาคำนวณเพื่อให้ได้ข้อมูลขั้นสุดท้าย ที่จะนำไปคำนวณทางสถิติ อย่างไรก็ตามก่อนที่จะได้มาซึ่งข้อมูล จะต้องมีความรู้ที่จะวิเคราะห์ก่อน การได้มาซึ่งตัวอย่างจึงเป็นสิ่งสำคัญที่ต้องคำนึงถึง เพื่อให้ได้ตัวอย่างที่เป็นตัวแทนที่ดีของประชากรที่เก็บตัวอย่างนั้นมา ผลวิเคราะห์ที่ได้จึงจะสามารถอธิบายลักษณะประชากรนั้นได้ใกล้เคียงถูกต้อง

จากที่กล่าวมาจะพบว่าตัวอย่างทางเคมีวิเคราะห์มีความแตกต่างจากตัวอย่างทางสถิติทั่วไปทางด้านอื่นที่ตัวอย่างจะแทนค่าข้อมูลที่จะนำมาคำนวณทางสถิติ หนึ่งตัวอย่างจึงให้หนึ่งข้อมูล เช่น นักเรียนชั้นมัธยมปลายที่สนใจเข้าศึกษาต่อมหาวิทยาลัยบูรพา ผู้สูงอายุที่ยังทำงานประจำ เป็นต้น แต่ตัวอย่างทางเคมีวิเคราะห์ต้องนำมาผ่านขั้นตอนการเตรียมตัวอย่างทางเคมี เช่น การสกัด การแยก การละลาย ก่อนที่จะนำไปตรวจวัดเพื่อให้ได้ข้อมูลดิบออกมา และผ่านการคำนวณเป็นข้อมูลขั้นสุดท้ายซึ่งเป็นค่าปริมาณสารที่ตรวจวัดได้ที่จะนำไปวิเคราะห์ด้วยสถิติ ดังนั้นทางเคมีวิเคราะห์ซึ่งมีการตรวจวัดซ้ำ (repeated measurement) ต่อหนึ่งตัวอย่างที่เตรียมมา จึงทำให้ได้ข้อมูลขั้นสุดท้ายมากขึ้น ตามจำนวนครั้งของการตรวจวัดซ้ำ ซึ่งก็คือขนาดตัวอย่าง (sample size; n) ที่เทียบได้กับจำนวนตัวอย่างในกลุ่มตัวอย่างที่ถูกสุ่มมาจากประชากรของทางสถิติที่ใช้ทั่วไป (Miller & Miller, 2010) โดยสรุป ขนาดตัวอย่าง (n) หมายถึง จำนวนครั้งที่ตรวจวัดซ้ำ (number of repeated measurement) ของตัวอย่างในงานวิจัยเชิงการทดลอง ตัวอย่างดังในตารางที่ 2.1 หรือจำนวนตัวอย่างที่ตรวจนับหรือสังเกต (number of observations) ในงานวิจัยเชิงสำรวจ ตัวอย่างดังในตารางที่ 2.2