

หนังสือ “การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรม SPSS” เล่มนี้ จัดพิมพ์เป็นครั้งที่ 3 มีการปรับปรุงแก้ไขจากการพิมพ์ในครั้งก่อนให้มีความสมบูรณ์มากยิ่งขึ้น โดยยังคงเนื้อหาสำคัญภายในเล่มซึ่งได้อธิบายเทคนิคทางสถิติที่มีความสำคัญต่อการพัฒนาผลิตภัณฑ์ ในอุตสาหกรรมเกษตร ประกอบด้วย กฎเกณฑ์ทางสถิติและการประยุกต์ใช้ในงานพัฒนาผลิตภัณฑ์ ตลอดจนการใช้โปรแกรมสำเร็จรูปทางสถิติ SPSS ในการวิเคราะห์ข้อมูลงานวิจัยและพัฒนา เริ่มตั้งแต่ การป้อนข้อมูล การเลือกใช้ค่าสิ่งสถิติเพื่อวิเคราะห์ การแปลผล และการเขียนรายงานผลการวิเคราะห์ทางสถิติ หนังสือเล่มนี้เหมาะสำหรับนิสิตนักศึกษาใช้ประกอบการเรียนการสอนในระดับมหาวิทยาลัยและบุคลากรที่ทำงานในด้านการวิจัยและพัฒนาผลิตภัณฑ์ทางด้านอุตสาหกรรมเกษตร

ISBN 978-616-586-105-2



ราคา 350 บาท



การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรม SPSS รศ.ดร.ปิณฑร ฤทธิเรืองเดช

การวิเคราะห์ข้อมูลงานวิจัยและ พัฒนาผลิตภัณฑ์โดยใช้โปรแกรม

SPSS

Data Analysis for Research and Product Development using SPSS

รองศาสตราจารย์ ดร.ปิณฑร ฤทธิเรืองเดช
ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร
มหาวิทยาลัยเกษตรศาสตร์

การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์ โดยใช้โปรแกรม SPSS

(พิมพ์ครั้งที่ 3 ฉบับปรับปรุง)

รองศาสตราจารย์ ดร. ปิติพร ฤทธิเรืองเดช

ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร

มหาวิทยาลัยเกษตรศาสตร์

การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรม SPSS

รองศาสตราจารย์ ดร.ปิติพร ฤทธิเรืองเดช

พิมพ์ครั้งที่ 1 พฤษภาคม 2561 จำนวน 250 เล่ม

พิมพ์ครั้งที่ 2 มิถุนายน 2561 จำนวน 250 เล่ม

พิมพ์ครั้งที่ 3 สิงหาคม 2564 จำนวน 250 เล่ม

(ฉบับปรับปรุง)

ราคา 350 บาท

สงวนลิขสิทธิ์ในประเทศไทยตาม พ.ร.บ.ลิขสิทธิ์

ห้ามผู้ใดพิมพ์ซ้ำ ลอกเลียน ส่วนหนึ่งส่วนใดของหนังสือเล่มนี้

นอกจากจะได้รับอนุญาตจากผู้เขียนเป็นลายลักษณ์อักษรเท่านั้น

ข้อมูลทางบรรณานุกรมของหอสมุดแห่งชาติ

ปิติพร ฤทธิเรืองเดช.

การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรม SPSS.

-- พิมพ์ครั้งที่ 3.-- กรุงเทพฯ : วิสด้า อินเทอร์เน็ต, 2564.

387 หน้า.

1. การออกแบบผลิตภัณฑ์ 2. เอสพีเอสเอส (โปรแกรมคอมพิวเตอร์). I. ชื่อเรื่อง.

658.57

ISBN 978-616-586-105-2

จัดพิมพ์และจัดจำหน่ายโดย :

รศ.ดร.ปิติพร ฤทธิเรืองเดช

ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเกษตรศาสตร์

50 ถนนงามวงศ์วาน แขวงลาดยาว เขตจตุจักร กรุงเทพฯ 10900

โทรศัพท์ 0-2562-5004 โทรสาร 0-2562-5005 Email: pitiporn@gmail.com

รูปเล่ม ปิติพร ฤทธิเรืองเดช

ออกแบบปก ณัฐกมล ไกรพจน์

ภาพถ่าย มินตรา นักรธรรม, สุภามาส หาญเพิ่มชัย, อริสรา หิริโอตัมปะ

พิสูจน์อักษร พชรนันท์ สุขแสงพนมรุ้ง, อริสรา หิริโอตัมปะ

โรงพิมพ์ บริษัท วิสด้า อินเทอร์เน็ต จำกัด โทร. 0-2015-4377 โทรสาร. 0-2015-4377

คำนำ

หนังสือเล่มนี้ได้จัดพิมพ์ขึ้นเป็นครั้งที่ 3 โดยได้มีการปรับปรุงแก้ไขจากการพิมพ์ในครั้งก่อนให้มีความสมบูรณ์มากยิ่งขึ้นเพื่อให้ผู้สนใจในงานวิจัยและพัฒนาผลิตภัณฑ์สามารถเลือกใช้สถิติและวิธีการในการวิเคราะห์ข้อมูลได้อย่างถูกต้องและเหมาะสม ซึ่งเหมาะสำหรับนิสิต นักศึกษา ครู อาจารย์ นักวิจัย และบุคลากรที่ทำงานในด้านพัฒนาผลิตภัณฑ์ทางอุตสาหกรรมเกษตร โดยผู้อ่านต้องมีพื้นฐานการเรียนรู้สถิติพื้นฐานมาก่อนบ้างแล้ว จุดมุ่งหมายของหนังสือเล่มนี้เพื่อให้ผู้อ่านสามารถเรียนรู้สถิติได้อย่างสนุกและเข้าใจ คิดวางแผนการทดลองเป็น สามารถนำสถิติไปประยุกต์ใช้ในงานวิจัยและนำเสนอผลการวิเคราะห์ได้อย่างถูกต้อง ซึ่งผู้เขียนได้เรียบเรียงจากหนังสือ ตำรา งานวิจัย และประสบการณ์ของผู้เขียนในการสอนและวิจัยทั้งในด้านหลักการพัฒนาผลิตภัณฑ์และการประยุกต์ใช้สถิติในงานพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร

เนื้อหาภายในเล่มประกอบด้วย ความสำคัญของงานวิจัยและพัฒนาผลิตภัณฑ์, สถิติพื้นฐานและการทบทวน, ความรู้พื้นฐานเกี่ยวกับการใช้โปรแกรมสำเร็จรูป SPSS for Windows, สถิติเชิงพรรณนาที่ใช้ในงานวิจัย, การทดสอบสมมติฐาน, การวิเคราะห์ผลทางสถิติด้วย t-test, การวิเคราะห์ความแปรปรวน, การวางแผนการทดลองแบบสุ่มสมบูรณ์, การวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ และการวางแผนแบบแฟกทอเรียล ซึ่งในแต่ละบทนอกจากจะอธิบายทฤษฎีทางสถิติและการประยุกต์ใช้ในงานพัฒนาผลิตภัณฑ์แล้วยังมีเนื้อหาเกี่ยวกับการใช้โปรแกรมสำเร็จรูปทางสถิติ SPSS ตั้งแต่การป้อนข้อมูล, การเลือกใช้คำสั่งสถิติเพื่อวิเคราะห์, การแปลความหมาย, การสรุปผล และการนำเสนอผลที่ได้จากการวิเคราะห์

ผู้เขียนหวังเป็นอย่างยิ่งว่าหนังสือเล่มนี้จะเป็นประโยชน์แก่ นิสิต นักศึกษา นักวิจัย และ บุคลากรที่ทำงานในด้านวิจัยและพัฒนาผลิตภัณฑ์ หากพบข้อผิดพลาดหรือมีข้อเสนอแนะประการใดในเนื้อหาของหนังสือเล่มนี้ กรุณาแจ้งผู้เขียนโดยตรงเพื่อจะได้ทำการปรับปรุงแก้ไขในการจัดพิมพ์ครั้งต่อไป

ท้ายนี้ผู้เขียนใคร่ขอน้อมรำลึกถึงพระคุณ บิดามารดา ครูอาจารย์ทุกท่านที่ได้ประสิทธิ์ประสาทความรู้ให้กับผู้เขียน และขอขอบคุณนิสิตทั้งระดับปริญญาตรีและปริญญาโทที่ผู้เขียนสอนและร่วมทำวิจัยซึ่งเปรียบเสมือนกระจกที่ส่องสะท้อนให้ผู้เขียนพัฒนาและปรับปรุงกระบวนการสอนและวิจัยให้ดียิ่งขึ้นต่อไปตลอดจนทุกท่านที่มีส่วนช่วยให้หนังสือเล่มนี้สำเร็จลุล่วงไปด้วยดี

รองศาสตราจารย์ ดร.ปิติพร ฤทธิเรืองเดช

สิงหาคม 2564

สารบัญ

หน้า

บทที่ 1 ความสำคัญของงานวิจัยและพัฒนาผลิตภัณฑ์	1
1.1 ความหมายและความสำคัญของงานพัฒนาผลิตภัณฑ์	1
1.2 สาเหตุที่ต้องมีการพัฒนาผลิตภัณฑ์ใหม่	3
1.3 ประเภทและระดับของผลิตภัณฑ์	10
1.4 ความหมายและประเภทของผลิตภัณฑ์ใหม่	13
1.5 กระบวนการพัฒนาผลิตภัณฑ์	17
1.6 บทบาทและความสำคัญของสถิติในงานวิจัยและพัฒนา	28
บทที่ 2 สถิติพื้นฐานและการทบทวน	32
2.1 ประเภทของสถิติสำหรับงานวิจัย	32
2.2 ประชากรและตัวอย่าง	34
2.3 ตัวแปร	35
2.4 ประเภทของข้อมูล	36
2.5 การทดลอง	42
2.6 ขั้นตอนในการวางแผนการทดลอง	42
2.7 สิ่งทดลอง	43
2.8 หน่วยทดลอง	44
2.9 ปัจจัยและระดับของปัจจัย	46
2.10 ความคลาดเคลื่อนในการทดลอง	48
2.11 การทำซ้ำ	48
2.12 การสุ่ม	49
2.13 ความสม่ำเสมอของหน่วยทดลอง	50
2.14 การบล็อก	50
บทที่ 3 ความรู้พื้นฐานเกี่ยวกับการใช้โปรแกรมสำเร็จรูป SPSS for Windows	51
3.1 การเริ่มต้นใช้งานโปรแกรม SPSS for Windows	51
3.2 ข้อกำหนดทั่วไปของโปรแกรม SPSS for Windows	52

	หน้า
3.3 การป้อนข้อมูลหรือนำข้อมูลเข้าสู่ Work sheet	60
3.4 เมนูสำหรับการวิเคราะห์ผลทางสถิติ	67
3.5 การบันทึกผลลัพธ์ (Output) ของโปรแกรม SPSS	69
บทที่ 4 สถิติเชิงพรรณนาที่ใช้ในงานวิจัย	72
4.1 ประเภทของสถิติเชิงพรรณนา	72
4.2 การเก็บรวบรวมข้อมูลเพื่อใช้ในการวิเคราะห์สถิติเชิงพรรณนา	73
4.3 สถิติเชิงพรรณนาสำหรับข้อมูลเชิงคุณภาพหรือเชิงกลุ่ม	78
4.4 สถิติเชิงพรรณนาสำหรับข้อมูลเชิงปริมาณ	93
4.5 การแจกแจงความถี่แบบ 2 ทางโดยใช้คำสั่ง Crosstabs	102
4.6 กรณีศึกษาการสร้างแบบสอบถาม การกำหนดตัวแปรและรหัสข้อมูลในโปรแกรม SPSS	105
บทที่ 5 การทดสอบสมมติฐาน	116
5.1 ความหมายของสมมติฐาน	116
5.2 ประเภทของสมมติฐาน	117
5.3 ประเภทของการทดสอบสมมติฐาน	118
5.4 ความคลาดเคลื่อนในการทดสอบสมมติฐานทางสถิติและระดับนัยสำคัญ	122
5.5 เขตวิกฤตและค่าวิกฤต	124
5.6 ขั้นตอนการทดสอบสมมติฐานทางสถิติ	126
5.7 หลักเกณฑ์ในการพิจารณาเลือกใช้สถิติในการทดสอบ	127
บทที่ 6 การวิเคราะห์ผลทางสถิติด้วย t-test	136
6.1 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 1 กลุ่มตัวอย่าง	136
6.1.1 การเขียนสมมติฐานทางสถิติเพื่อทดสอบค่าเฉลี่ย 1 กลุ่มตัวอย่าง	137
6.1.2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ค่าเฉลี่ยด้วยวิธี One-sample t-test	138
6.1.3 การสรุปผลลัพธ์การวิเคราะห์ค่าเฉลี่ยด้วยวิธี One-sample t-test	141
6.1.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี One-sample t-test	144
6.2 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง	145
6.2.1 สถิติ t-test ที่ใช้ทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง	145

	หน้า
6.2.2 การเขียนสมมติฐานทางสถิติสำหรับทดสอบค่าเฉลี่ย 2 กลุ่มตัวอย่าง	146
6.3 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่างที่ไม่เป็นอิสระกัน	150
6.3.1 ลักษณะของข้อมูลกลุ่มตัวอย่างที่ไม่เป็นอิสระกัน	150
6.3.2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test	151
6.3.3 การสรุปผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test	156
6.3.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test	157
6.3.5 กรณีศึกษาการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test	158
6.4 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่างที่เป็นอิสระกัน	163
6.4.1 การคำนวณค่าสถิติ t ด้วยวิธี Independent-samples t-test	163
6.4.2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test	164
6.4.3 การสรุปผลลัพธ์การวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test	170
6.4.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test	172
6.4.5 กรณีศึกษาการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test	173
บทที่ 7 การวิเคราะห์ความแปรปรวน	179
7.1 การวิเคราะห์ความแปรปรวน	180
7.2 ชนิดของความแปรปรวน	181
7.3 ประเภทของการวิเคราะห์ความแปรปรวน	182
7.4 ขั้นตอนการวิเคราะห์ความแปรปรวน	183
7.5 การวิเคราะห์ความแปรปรวนแบบทางเดียว	186
7.6 คำสั่ง SPSS ในการวิเคราะห์ความแปรปรวนแบบทางเดียว	187
7.7 กรณีศึกษาการวิเคราะห์ความแปรปรวนแบบทางเดียว	198
7.8 การวิเคราะห์ความแปรปรวนแบบสองทาง	208
7.9 คำสั่ง SPSS ในการวิเคราะห์ความแปรปรวนแบบสองทาง	211
7.10 กรณีศึกษาการวิเคราะห์ความแปรปรวนแบบสองทาง	229

	หน้า
บทที่ 8 การวางแผนการทดลองแบบสุ่มสมบูรณ์	243
8.1 ข้อดีและข้อเสียของการวางแผนการทดลองแบบสุ่มสมบูรณ์	244
8.2 ตัวแบบ	244
8.3 การตั้งสมมติฐานเพื่อการทดสอบความสัมพันธ์หรือทดสอบเปรียบเทียบค่าเฉลี่ย	246
8.4 ตัวอย่างรูปแบบข้อมูลจากการวางแผนการทดลองแบบสุ่มสมบูรณ์	247
8.5 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD	249
8.5.1 ตัวอย่างข้อมูล	249
8.5.2 การเขียนสมมติฐานเพื่อทดสอบ	249
8.5.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD	250
8.5.4 การอ่านผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ CRD	255
8.5.5 การนำเสนอผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ CRD	258
8.6 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD	259
บทที่ 9 การวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์	271
9.1 ข้อดีและข้อเสียของการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์	272
9.2 การบล็อก	272
9.3 ตัวแบบ	273
9.4 การตั้งสมมติฐานเพื่อการทดสอบความสัมพันธ์หรือทดสอบเปรียบเทียบค่าเฉลี่ย	275
9.5 ตัวอย่างรูปแบบข้อมูลจากการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์	276
9.6 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD	278
9.6.1 ตัวอย่างข้อมูล	279
9.6.2 การเขียนสมมติฐานเพื่อทดสอบ	279
9.6.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD	280
9.6.4 การอ่านผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ RCBD	288
9.6.5 การนำเสนอผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ RCBD	290
9.7 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD	291

บทที่ 10 การทดลองแบบแฟกทอเรียล	301
10.1 การทดลองแบบแฟกทอเรียล	301
10.2 การจัดกลุ่มสิ่งทดลอง	302
10.3 นิยามของอิทธิพลหลักและอิทธิพลร่วมระหว่างปัจจัย	308
10.4 การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย	310
10.4.1 ความแปรปรวนของการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย	310
10.4.2 การเขียนสมมติฐานทางสถิติสำหรับการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย	312
10.4.3 การทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ	313
10.4.4 คำสั่ง SPSS ในการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย	313
10.5 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 3×2 Factorial in CRD	329
10.6 การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย	343
10.6.1 ความแปรปรวนของการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย	344
10.6.2 การเขียนสมมติฐานทางสถิติสำหรับการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย	346
10.6.3 การทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ	347
10.6.4 คำสั่ง SPSS ในการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย	347
10.7 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 2 ² Factorial in RCBD	366
บรรณานุกรม	382
ดรชนี	384



บทที่ 1

ความสำคัญของงานวิจัยและ พัฒนาผลิตภัณฑ์

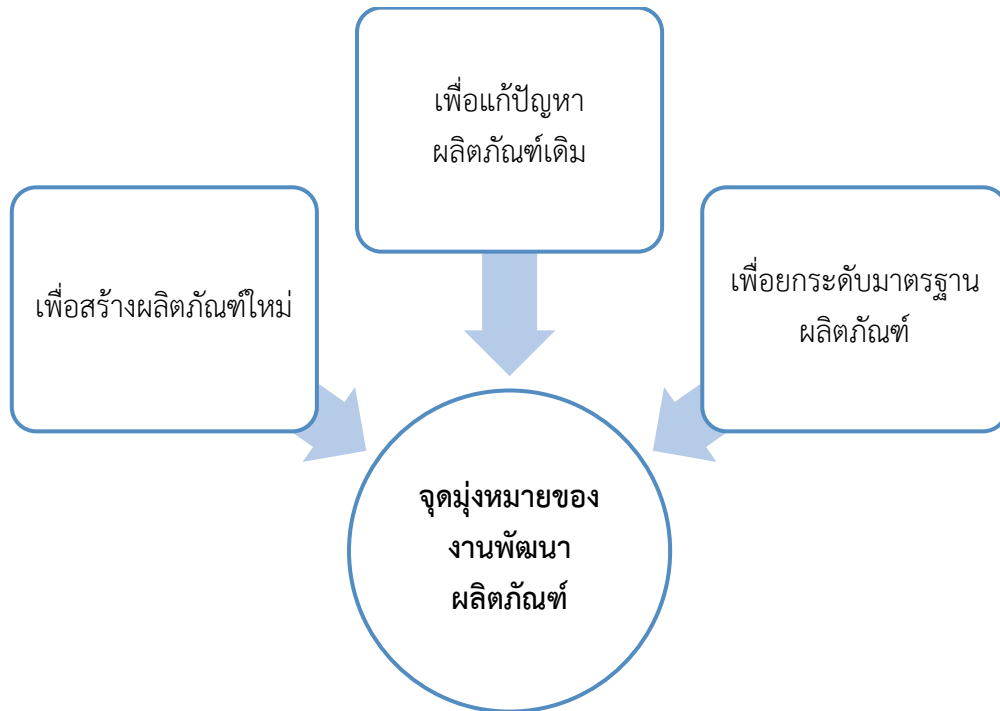
ประเทศไทยเป็นหนึ่งในหลายประเทศของโลกที่มีอาหารบริโภคอย่างเพียงพอและยังมีเหลือเพื่อการส่งออกทั้งในรูปแบบของ สินค้าเกษตร สินค้าเกษตรแปรรูป และอาหาร การแปรรูปอุตสาหกรรมเกษตร คือ การนำวัตถุดิบจากการผลิตทางการเกษตรและวัตถุดิบที่ได้จากธรรมชาติจำนวนมากมาดำเนินการแปรรูปด้วยเครื่องจักรกลอย่างต่อเนื่อง เพื่อให้ได้ปริมาณของผลิตภัณฑ์ที่มีคุณภาพ และมีต้นทุนการผลิตตามต้องการภายในระยะเวลาที่กำหนด ผลที่ได้จากการดำเนินการนี้ คือ 1) ผลิตภัณฑ์ที่ต้องการ (Product) 2) ผลพลอยได้ (By product) และ 3) ของเหลือ (Waste) การหาแนวทางการใช้ประโยชน์จากทั้งสามส่วนสามารถทำได้โดยอาศัยหลักการพัฒนาผลิตภัณฑ์ทางอุตสาหกรรมเกษตร ซึ่งจุดมุ่งหมายสำคัญของการพัฒนาผลิตภัณฑ์ทางอุตสาหกรรมเกษตร คือ การเพิ่มมูลค่าให้ผลผลิตทางการเกษตรโดยยกระดับคุณภาพผลิตภัณฑ์ให้ได้ตามมาตรฐาน และเป็นผลิตภัณฑ์ที่สามารถตอบสนองต่อความต้องการของผู้บริโภคได้

1.1 ความหมายและความสำคัญของงานพัฒนาผลิตภัณฑ์

งานวิจัยและพัฒนา (Research and Development) มีความสำคัญต่ออุตสาหกรรมทุกแขนง โดยเฉพาะงานพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร ประเทศไทยมีวัตถุดิบเกษตรมากมายชนิดและเกือบทุกชนิดมีศักยภาพการผลิตสูงแต่ไม่อาจเก็บรักษาได้นาน มีการเสื่อมคุณภาพเน่าเสียได้ง่าย และยังมีมูลค่าทางเศรษฐกิจต่ำ ดังนั้นการนำหลักการพัฒนาผลิตภัณฑ์มาประยุกต์ใช้เพื่อหาแนวทางการแปรรูปวัตถุดิบเกษตรให้เป็นผลิตภัณฑ์นานาชาติที่ตลาดต้องการ และการยกระดับคุณภาพของผลิตภัณฑ์เดิมที่มีอยู่แล้วให้ได้มาตรฐานจึงเป็นเรื่องที่อุตสาหกรรมให้ความสำคัญเป็นอย่างยิ่ง

การพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร หมายถึง การดำเนินงานพัฒนาการใช้ประโยชน์จากผลผลิตทางการเกษตรทุกชนิดทั้งที่เป็น อาหาร กึ่งอาหาร และ ไม่ใช่อาหาร ที่มีมูลค่าต่ำ แล้วทำการวิจัยหาแนวทางเพื่อแปรรูปเป็นผลิตภัณฑ์สำเร็จรูปหรือผลิตภัณฑ์สำเร็จรูปที่มีมูลค่าสูงกว่า และสินค้าที่พัฒนาแล้วเป็นที่ต้องการของตลาด ตลอดจนรวมไปถึงการกำหนดควบคุมคุณภาพวัตถุดิบและผลิตภัณฑ์ และการยกระดับมาตรฐานของผลิตภัณฑ์

จุดมุ่งหมายหลักของงานพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร คือ เพื่อสร้างผลิตภัณฑ์ใหม่ เพื่อแก้ปัญหาผลิตภัณฑ์เดิมให้คงอยู่ในตลาดโดยการปรับปรุงคุณภาพผลิตภัณฑ์ให้ตอบสนองต่อความต้องการของผู้บริโภค และเพื่อยกระดับคุณภาพผลิตภัณฑ์และบรรจุภัณฑ์ให้ตอบสนองต่อความต้องการของผู้บริโภค ดังแสดงในภาพที่ 1.1



ภาพที่ 1.1 จุดมุ่งหมายของงานพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร

กระบวนการพัฒนาผลิตภัณฑ์เป็นขั้นตอนสำคัญของอุตสาหกรรมเกษตร เป็นการทำงานร่วมกันอย่างมีระบบแบบแผนของหลายหน่วยงานที่มีหน้าที่รับผิดชอบต่างกัน ลักษณะการทำงานของหน่วยงานพัฒนาผลิตภัณฑ์เป็นการทำงานที่เรียกว่า “วิจัยอุตสาหกรรมประยุกต์ (Applied Industrial Research)” เนื่องจากไม่ได้มีเป้าหมายแค่ผลิตภัณฑ์สุดท้ายที่พัฒนาได้ แต่ยังรวมถึงงานวิจัยทางด้านวัตถุดิบ เทคโนโลยีด้านการผลิต เทคโนโลยีด้านการบรรจุ การพัฒนาสูตร การควบคุมคุณภาพ การประกันคุณภาพ การจัดการ และการตลาด โดยเป็นการใช้องค์ความรู้ในด้านดังกล่าวร่วมกันเพื่อพัฒนาผลิตภัณฑ์ ทั้งนี้การประยุกต์ใช้องค์ความรู้เหล่านี้ต้องอาศัยหลักการด้านสถิติเพื่อช่วยในการวางแผนงานวิจัยและตัดสินใจ ดังนั้นจึงกล่าวได้ว่าสถิติเป็นเครื่องมือที่สำคัญสำหรับงานวิจัย การที่นักวิจัยจะทำวิจัยให้ประสบความสำเร็จได้นั้นต้องมีความรู้ความเข้าใจในศาสตร์ของงานวิจัยนั้นควบคู่ไปกับความรู้ทางสถิติ เช่น การวิจัยหาสภาวะที่เหมาะสมในกระบวนการผลิต นอกจากนักวิจัยจะต้องมีความรู้ด้านวัตถุดิบ กระบวนการผลิต และการวัดค่าคุณภาพแล้วยังต้องมีความรู้ทางสถิติด้วย

หน้าที่หลักของงานวิจัยและพัฒนาผลิตภัณฑ์ มีดังนี้

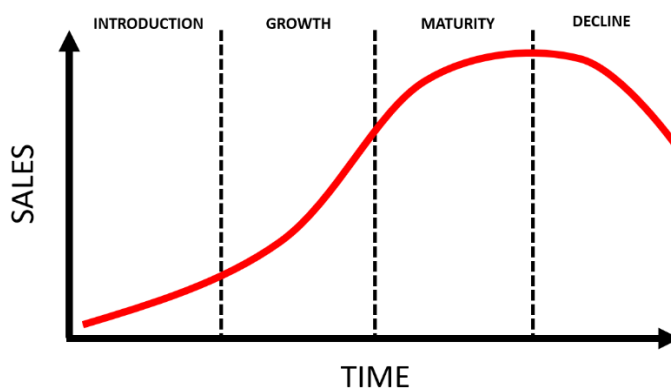
- รวบรวมความรู้ทางวิทยาศาสตร์ เทคโนโลยี และการตลาดที่ทันสมัยที่จะก่อให้เกิดประโยชน์ต่องานส่วนรวมและการนำไปใช้ประโยชน์
- ใช้ความรู้ทางวิทยาศาสตร์ เทคโนโลยี และการตลาด เพื่อพัฒนาผลิตภัณฑ์ใหม่ๆ ตลอดจนปรับปรุงคุณภาพผลิตภัณฑ์และกระบวนการผลิตให้สัมพันธ์กัน
- การควบคุมคุณภาพผลิตภัณฑ์ การกำหนดมาตรฐานและคุณลักษณะวัตถุที่ต้องการ
- ให้ความร่วมมือที่เกี่ยวกับการใช้วิทยาศาสตร์ เทคโนโลยี และ การตลาด กับแขนงอื่นๆ

1.2 สาเหตุที่ต้องมีการพัฒนาผลิตภัณฑ์ใหม่

สาเหตุที่ต้องมีการพัฒนาผลิตภัณฑ์ใหม่ มีดังนี้

(1) ผลิตภัณฑ์มีวงจรชีวิต (Product life cycle) สิ้นลง

วงจรชีวิตผลิตภัณฑ์แสดงถึงการเปลี่ยนแปลงยอดขายของผลิตภัณฑ์ตั้งแต่เริ่มนำผลิตภัณฑ์ออกสู่ตลาดจนกระทั่งผลิตภัณฑ์มียอดขายตกต่ำจนต้องนำออกจากตลาด ผลิตภัณฑ์แต่ละตัวจะมีวงจรชีวิตผลิตภัณฑ์ที่แตกต่างกัน เมื่อผลิตภัณฑ์ได้รับการยอมรับมียอดขายสูงจะทำให้เกิดการแข่งขันจากผลิตภัณฑ์ของบริษัทคู่แข่งมาแย่งส่วนแบ่งตลาดทำให้ยอดขายลดลง หากผลิตภัณฑ์เดิมสูญเสียผลิตภัณฑ์ของบริษัทคู่แข่งไม่ได้ หรือผลิตภัณฑ์มีความล้าสมัยจนไม่สามารถตอบสนองความต้องการของผู้บริโภคได้ ในที่สุดผลิตภัณฑ์นั้นก็จะหายไปจากตลาดแล้วผลิตภัณฑ์ใหม่ก็จะเข้ามาแทนที่เป็นวัฏจักรเรื่อยไป ผลิตภัณฑ์ในปัจจุบันมีวงจรชีวิตที่สั้นลงอันเนื่องมาจากความเจริญก้าวหน้าของเทคโนโลยีที่ทำให้บริษัทคู่แข่งสามารถผลิตเลียนแบบได้ง่าย อีกทั้งพฤติกรรมของผู้บริโภคที่เปลี่ยนแปลงอย่างรวดเร็วจึงมีผลทำให้ความภักดีต่อตราสินค้าลดลง ภาพที่ 1.2 แสดงกราฟวงจรชีวิตของผลิตภัณฑ์ ซึ่งแบ่งออกได้เป็น 4 ช่วง ได้แก่ ช่วงแนะนำผลิตภัณฑ์ (Introduction stage) ช่วงเจริญเติบโต (Growth stage) ช่วงอิ่มตัว (Maturity stage) และ ช่วงถดถอย (Decline stage)



ภาพที่ 1.2 กราฟวงจรชีวิตผลิตภัณฑ์ (Product life cycle)

วงจรชีวิตผลิตภัณฑ์ แบ่งออกได้เป็น 4 ช่วง ดังนี้

ช่วงที่ 1 ช่วงแนะนำผลิตภัณฑ์ (Introduction stage) เป็นช่วงเริ่มนำผลิตภัณฑ์ออกวางจำหน่ายในท้องตลาด ช่วงนี้ยอดขายของผลิตภัณฑ์ยังไม่สูงเนื่องจากผู้บริโภคยังไม่รู้จักผลิตภัณฑ์ จึงต้องมุ่งเน้นการส่งเสริมการตลาดและการสื่อสารเพื่อนำเสนอรายละเอียดของผลิตภัณฑ์ให้ผู้บริโภคทราบ อาจใช้สื่อโฆษณาประเภทต่างๆ ที่สามารถเข้าสู่กลุ่มเป้าหมายได้ ใช้พนักงานขายที่มีความกระตือรือร้น การแจกผลิตภัณฑ์ตัวอย่างหรือการสาธิตเพื่อให้ผู้บริโภคได้ลองใช้หรือบริโภคผลิตภัณฑ์ ช่วงนี้เป็นช่วงที่ต้องใช้เงินลงทุนจำนวนมากในการสื่อสารให้ผู้บริโภคทราบรายละเอียดของผลิตภัณฑ์ และต้องดำเนินการอย่างรวดเร็วก่อนที่บริษัทคู่แข่งจะเริ่มนำผลิตภัณฑ์ออกสู่ตลาด

ช่วงที่ 2 ช่วงเติบโต (Growth stage) ช่วงนี้ยอดขายของผลิตภัณฑ์เพิ่มขึ้นอย่างรวดเร็ว เนื่องจากผู้บริโภคยอมรับผลิตภัณฑ์และรู้จักผลิตภัณฑ์มากขึ้น บริษัทจะได้รับผลประโยชน์จากกำไรที่ค่อนข้างสูง ช่วงนี้จะเริ่มมีคู่แข่งทางการตลาดเกิดขึ้นทำให้ผู้บริโภคเริ่มเปรียบเทียบผลิตภัณฑ์ของบริษัทกับของบริษัทคู่แข่ง มีผลทำให้ส่วนแบ่งตลาดลดลง ดังนั้นบริษัทต้องมุ่งเน้นการส่งเสริมการขายที่ทำให้ผู้บริโภคภักดีต่อตราสินค้าของบริษัท ต้องเจาะจงซื้อผลิตภัณฑ์ของบริษัทแทนที่จะซื้อของคู่แข่ง เปลี่ยนการส่งเสริมการตลาดจากการรับรู้เป็นการภักดีต่อตราสินค้า นอกจากนี้จากการที่ยอดขายผลิตภัณฑ์เพิ่มขึ้นบริษัทจึงต้องเพิ่มในเรื่องการผลิตและการกระจายสินค้าให้ทั่วถึงเพื่อไม่ให้เกิดปัญหาสินค้าขาดตลาด

ช่วงที่ 3 ช่วงอิมมัตว (Maturity stage) ช่วงนี้ยอดขายผลิตภัณฑ์ของบริษัทเริ่มลดลงอันเป็นผลเนื่องมาจากมีผลิตภัณฑ์ของบริษัทคู่แข่งมาแย่งส่วนแบ่งตลาดมากขึ้น ผู้บริโภคมีโอกาสเลือกซื้อผลิตภัณฑ์ชนิดเดียวกันจากหลายบริษัทมากขึ้น ประกอบกับผู้บริโภคมีความสนใจในตัวผลิตภัณฑ์ลดลง จึงมีผลทำให้ยอดขายไม่เพิ่มขึ้นมากเหมือนช่วงเติบโต อัตราการเพิ่มค่อนข้างหยุดนิ่ง ดังนั้นในช่วงนี้บริษัทต้องมุ่งเน้นการแย่งส่วนแบ่งตลาดเดิมให้มากที่สุด หรือการเข้าสู่ตลาดใหม่โดยดึงดูดผู้บริโภคที่ไม่เคยใช้ผลิตภัณฑ์ให้หันมาสนใจในตัวผลิตภัณฑ์ ต้องมีการปรับปรุงเปลี่ยนแปลงผลิตภัณฑ์ให้แปลกใหม่เพื่อสร้างความแตกต่างของผลิตภัณฑ์ให้เหนือคู่แข่ง (Product differentiation) เช่น ปรับปรุงบรรจุภัณฑ์ให้ทันสมัย รวมทั้งการปรับปรุงส่วนประสมทางการตลาด เช่น การลดราคาสินค้า เพิ่มช่องทางการจำหน่ายออนไลน์

ช่วงที่ 4 ช่วงถดถอย (Decline stage) ช่วงนี้ยอดขายเริ่มลดลงเรื่อยๆ อาจเนื่องมาจากมีผลิตภัณฑ์คู่แข่งเกิดขึ้นมาก ส่วนแบ่งตลาดจึงลดลงจนต้องนำผลิตภัณฑ์ออกจากตลาดเพื่อไม่ให้บริษัทขาดทุน ช่วงนี้บริษัทต้องใช้กลยุทธ์ลดราคาสินค้าให้มากเพื่อเร่งระบายผลิตภัณฑ์ที่มีอยู่ออกจากตลาดให้มากที่สุด เนื่องจากผู้บริโภคไม่ได้ทำการเลิกใช้ผลิตภัณฑ์ในทันที ดังนั้นบริษัทยังคงจำหน่ายผลิตภัณฑ์ได้อยู่ หากบริษัทตระหนักว่าการคงอยู่ของผลิตภัณฑ์ในท้องตลาดทำให้เกิดการขาดทุนของบริษัท ดังนั้นบริษัทควรหยุดการผลิตผลิตภัณฑ์นั้นแล้วทำการพัฒนาผลิตภัณฑ์ใหม่เพื่อตอบสนองความต้องการของผู้บริโภคต่อไป

(2) คู่แข่งมากขึ้น

จากที่กล่าวมาในข้อ (1) เมื่อยอดขายของผลิตภัณฑ์ในตลาดเพิ่มขึ้นจึงก่อให้เกิดคู่แข่งทางการตลาดมากขึ้น ประกอบกับความก้าวหน้าทางด้านเทคโนโลยีที่มีกำลังผลิตสูงมีผลทำให้การแข่งขันเป็นไปอย่างรวดเร็ว จึงมีผลทำให้ยอดขายของผลิตภัณฑ์ลดลง ดังนั้นจึงเป็นเรื่องที่หลีกเลี่ยงไม่ได้ที่นักพัฒนาผลิตภัณฑ์ต้องเข้ามามีบทบาทในแต่ละช่วงของวงจรชีวิตผลิตภัณฑ์ การรวบรวมข้อมูลทางการตลาดและผลิตภัณฑ์ในท้องตลาดจะทำให้ทราบถึงสิ่งที่ผู้บริโภคต้องการ สามารถนำข้อมูลดังกล่าวมาใช้ในการพัฒนากลยุทธ์ทั้งทางการตลาด เช่น การปรับราคาสินค้า และกลยุทธ์ทางด้านตัวผลิตภัณฑ์ เช่น การสร้างความแตกต่างของผลิตภัณฑ์ให้เหนือคู่แข่ง

ยกตัวอย่างเช่น ผลิตภัณฑ์ชาเขียวพร้อมดื่มซึ่งมีการแข่งขันกันอย่างสูงระหว่างผู้นำทางการตลาดสองบริษัท ทั้งสองบริษัทได้พัฒนาผลิตภัณฑ์ชาเขียวพร้อมดื่มให้มีรสชาติหลากหลายมากยิ่งขึ้นเพื่อเพิ่มทางเลือกให้แก่ผู้บริโภค ซึ่งพบว่ามี การเพิ่มรสชาติของผลิตภัณฑ์หลายรสชาติ เช่น ชาเขียวรสน้ำผึ้งผสมมะนาว ชาเขียวผสมเก็กฮวย ชาเขียวผสมจุกข้าวญี่ปุ่น ชาเขียวผสมน้ำนมข้าวโพด



ภาพที่ 1.3 ตัวอย่างผลิตภัณฑ์ชาเขียวพร้อมดื่มที่จำหน่ายในท้องตลาด

(3) ผลิตภัณฑ์มีความซับซ้อนมากขึ้น

ผลิตภัณฑ์มีความซับซ้อนมากขึ้นอันเนื่องมาจากการพัฒนาทางด้านเทคโนโลยีใหม่ๆ เช่น ผลิตภัณฑ์ลูกอมสอดไส้ช็อกโกแลต ผลิตภัณฑ์โยเกิร์ตที่มีส่วนผสมของโพรไบโอติก ผลิตภัณฑ์น้ำผลไม้สกัดเย็น (Cold pressed HPP) ผลิตภัณฑ์ข้าวหุงสุกบรรจุกระป๋อง และ ผลิตภัณฑ์เครื่องดื่มเกลือแร่พร้อมดื่ม ผลิตภัณฑ์เหล่านี้ได้จากการนำเทคโนโลยีด้านการแปรรูปที่ทันสมัยเข้ามาช่วยในการผลิตเพื่อตอบสนองความต้องการของผู้บริโภคในด้านความสะดวกในการขนส่งและการบริโภค รวมทั้งด้านคุณค่าทางโภชนาการ โดยปัจจุบันมุ่งเน้นพัฒนาผลิตภัณฑ์อาหารที่ผ่านการแปรรูปขั้นต่ำ (Minimally processed foods) มากขึ้น ดังนั้นนักพัฒนาผลิตภัณฑ์ต้องคอยติดตามข้อมูลข่าวสารทางด้านเทคโนโลยีอยู่เสมอเพื่อนำมาใช้ในการพัฒนาผลิตภัณฑ์ของบริษัทให้เป็นที่ต้องการของผู้บริโภคที่มีพฤติกรรมบริโภคเปลี่ยนแปลงไปอย่างรวดเร็ว



ภาพที่ 1.4 ตัวอย่างผลิตภัณฑ์ที่มีการนำเทคโนโลยีที่ทันสมัยมาใช้ในการกระบวนการผลิต

(4) การพัฒนาทางด้านบรรจุภัณฑ์

ในปัจจุบันมีความต้องการผลิตภัณฑ์อาหารที่ผ่านการแปรรูปขั้นต่ำ (Minimally processed foods) มากขึ้น ประกอบกับความเจริญก้าวหน้าทางด้านโลจิสติกส์ (Logistics) เช่น การกระจายสินค้าทางเครื่องบิน ทางระบบขนส่งของไปรษณีย์และบริษัทเอกชน จึงทำให้การพัฒนาบรรจุภัณฑ์เข้ามามีส่วนสำคัญอย่างมากในการพัฒนาผลิตภัณฑ์ใหม่เพื่อสนองความต้องการของผู้บริโภค ในปัจจุบันมีการพัฒนาบรรจุภัณฑ์ในหลายรูปแบบเพื่อเพิ่มความสะดวกสบายในแง่การใช้สอย รวมทั้งการปกป้องคุณภาพของผลิตภัณฑ์ การชะลอการเสื่อมเสียที่อาจเกิดขึ้น การบ่งชี้คุณภาพของผลิตภัณฑ์ และการใช้วัสดุที่ย่อยสลายได้ในการผลิตบรรจุภัณฑ์ที่เป็นมิตรกับสิ่งแวดล้อม

โดยทั่วไปหน้าที่หลักของบรรจุภัณฑ์ มี 4 ด้าน ได้แก่ 1) ปกป้อง (Protection) การเสื่อมเสียจากสภาพแวดล้อมภายนอก, 2) สื่อสาร (Communication) รายละเอียดของผลิตภัณฑ์ให้ผู้บริโภคทราบ, 3) ก่อให้เกิดความสะดวก (Convenience) ต่อการขนส่งและการใช้งาน และ 4) บรรจุ (Containment) ผลิตภัณฑ์ในรูปแบบและปริมาณต่างๆ โมเดลหน้าที่หลักของบรรจุภัณฑ์แสดงดังในภาพที่ 1.5 ในปัจจุบันได้เกิดนวัตกรรมทางด้านบรรจุภัณฑ์ที่นอกจากจะตอบสนองหน้าที่ทั้งสี่ด้านดังกล่าวแล้วยังมีส่วนเพิ่มเติมที่สามารถตอบสนองต่อความต้องการของผู้บริโภคในด้านอื่นได้ด้วย เช่น บรรจุภัณฑ์แบบแอคทีฟ (Active packaging) ทำหน้าที่ปกป้องคุณภาพของผลิตภัณฑ์ทั้งทางด้านกายภาพ เคมี และประสาทสัมผัส รวมทั้งด้านความปลอดภัย ทำให้ผลิตภัณฑ์มีอายุการเก็บรักษาที่นานขึ้น เช่น มีสารต้านการเกิดปฏิกิริยาออกซิเดชัน, บรรจุภัณฑ์อินเทลลิเจนท์ (Intelligent packaging) ทำหน้าที่ให้ข้อมูลหรือแสดงให้ผู้บริโภคทราบว่าคุณภาพของผลิตภัณฑ์ในภาชนะบรรจุนั้นเป็นอย่างไร เช่น มีตัวชี้วัดระดับความสุกหรือการเน่าเสียของผลิตภัณฑ์ได้ และบรรจุภัณฑ์ฉลาด (Smart packaging) คือ บรรจุภัณฑ์ที่ทำหน้าที่ได้ทั้งแบบบรรจุภัณฑ์แอคทีฟและบรรจุภัณฑ์อินเทลลิเจนท์



ภาพที่ 1.5 โมเดลหน้าที่หลักของบรรจุภัณฑ์
ที่มา : Yam et al. (2005)



ภาพที่ 1.6 ตัวอย่างบรรจุภัณฑ์ที่เป็นมิตรกับสิ่งแวดล้อม

(5) การเพิ่มชนิดของผลิตภัณฑ์สำหรับผู้บริโภค

พฤติกรรมและความต้องการของผู้บริโภคมีความหลากหลายขึ้นอยู่กับหลายปัจจัย เช่น เพศ อายุ ระดับการศึกษา อาชีพ ระดับเงินเดือน ถิ่นที่อยู่อาศัย ดังนั้นการตอบสนองความต้องการของผู้บริโภคที่มีความหลากหลายคือการเพิ่มชนิดของผลิตภัณฑ์เพื่อเป็นทางเลือกให้แก่ผู้บริโภคที่มีความต้องการแตกต่างกัน ตัวอย่างเช่น การพัฒนาผลิตภัณฑ์แกงเขียวหวานในปัจจุบัน ซึ่งมีตั้งแต่เครื่องแกงสำเร็จรูปบรรจุถุง เครื่องแกงสำเร็จรูปชนิดผง น้ำแกงเขียวหวานบรรจุกล่อง และผลิตภัณฑ์จากเครื่องแกงเขียวหวาน เช่น บะหมี่กึ่งสำเร็จรูปรสแกงเขียวหวาน ปลากระป๋องรสแกงเขียวหวาน



ภาพที่ 1.7 ตัวอย่างผลิตภัณฑ์จากแกงเขียวหวาน

(6) การลดความจำกัดทางด้านฤดูกาลของผลิตภัณฑ์

เนื่องจากผลผลิตทางการเกษตรบางอย่างมีตามฤดูกาลและสามารถเพาะปลูกได้ดีเฉพาะในบางพื้นที่ ประกอบกับผลผลิตทางการเกษตรมักเสื่อมเสียจากสภาพแวดล้อมได้ง่าย เช่น สภาพอากาศ การขนส่ง จึงทำให้เกิดข้อจำกัดในการใช้ประโยชน์ ดังนั้นการพัฒนาผลิตภัณฑ์จากวัตถุดิบทางการเกษตรนอกจากจะลดข้อจำกัดทางด้านฤดูกาลทำให้สามารถมีผลิตภัณฑ์ไว้บริโภคได้ตลอดปีแล้ว ยังเป็นการลดการสูญเสียที่เกิดขึ้นในระหว่างการเก็บรักษาได้อีกด้วย ตัวอย่างเช่น การนำสตอร์วเบอร์รี่สดมาแปรรูปได้หลายรูปแบบเพื่อนำไปใช้ทั้งบริโภคโดยตรงและนำไปใช้เป็นวัตถุดิบสำหรับผลิตผลิตภัณฑ์ชนิดอื่นต่อไป เช่น แยมสตอร์วเบอร์รี่ น้ำสตอร์วเบอร์รี่เข้มข้น สตอร์วเบอร์รี่อบแห้ง โยเกิร์ตใส่สตอร์วเบอร์รี่ น้ำเชื่อมรสสตอร์วเบอร์รี่ ดังนั้นสิ่งสำคัญที่นักพัฒนาผลิตภัณฑ์ต้องคำนึงถึง คือ วัตถุดิบที่ต้องใช้มีเพียงพอตลอดระยะเวลาการผลิตหรือไม่ และหากมีเพียงบางฤดูกาลนักพัฒนาผลิตภัณฑ์จะเก็บรักษาวัตถุดิบในรูปแบบใด



ภาพที่ 1.8 ตัวอย่างผลิตภัณฑ์จากสตอร์วเบอร์รี่

(7) ความต้องการผลิตภัณฑ์จากต่างประเทศ

ความต้องการผลิตภัณฑ์จากต่างประเทศ เช่น บะหมี่กึ่งสำเร็จรูป ขนมอบเคี้ยว ซอสปรุงรส เครื่องสำอาง ผลิตภัณฑ์อาหารสัตว์ พบว่ามีปริมาณที่ค่อนข้างสูง ดังนั้นการพัฒนาผลิตภัณฑ์ดังกล่าวโดยใช้วัตถุดิบที่มีอยู่ในประเทศ นอกจากจะช่วยเพิ่มมูลค่าให้ผลผลิตทางการเกษตรแล้วยังช่วยลดการนำเข้าอีกทางหนึ่งด้วย



ภาพที่ 1.9 ตัวอย่างผลิตภัณฑ์บะหมี่กึ่งสำเร็จรูปที่นำเข้าจากต่างประเทศและที่ผลิตในประเทศ

(8) ขนาดของร้านค้าปลีกใหญ่ขึ้น

ในปัจจุบันมีการแข่งขันของร้านค้าปลีกอย่างรุนแรง ร้านค้าปลีกมีขนาดใหญ่ขึ้นจึงเพิ่มพื้นที่รองรับผลิตภัณฑ์ได้มากขึ้น ดังนั้นในการพัฒนาผลิตภัณฑ์ต้องทำให้ผลิตภัณฑ์ของบริษัทมีจุดเด่นสามารถจำได้ง่าย และไม่ทำให้ผู้บริโภคสับสน



ภาพที่ 1.10 ตัวอย่างการจัดวางสินค้าในร้านค้าปลีก

(9) ข้อกำหนดของกฎหมายที่เปลี่ยนไป

การเปลี่ยนแปลงข้อกำหนดกฎหมายบางประการ ทำให้ผู้ผลิตต้องเปลี่ยนแปลงปรับสูตรผลิตภัณฑ์ให้เป็นไปตามนโยบายนั้น เช่น กรมสรรพสามิตได้ประกาศอัตราภาษีใหม่เมื่อเดือนกันยายน 2560 หรือที่เรียกว่า ภาษีน้ำตาล เพื่อให้ผู้ผลิตและผู้ประกอบการลดปริมาณน้ำตาลที่ผสมลงไปเครื่องดื่ม และเพื่อปรับเปลี่ยนพฤติกรรมของผู้บริโภคให้ลดการบริโภคเครื่องดื่มที่มีการใส่น้ำตาลเข้าไปมากๆ ทำให้ผู้ผลิตเครื่องดื่มในประเทศไทยต้องปรับสูตรและปรับกลยุทธ์ทางการตลาดกระตุ้นการตัดสินใจซื้อเพื่อเพิ่มยอดขาย โดยออกกลยุทธ์ไปในทิศทางเดียวกัน คือ นำเสนอเครื่องดื่มใหม่ๆ เพิ่มเติมโดยเน้นจุดขายสูตรน้ำตาลน้อยเพื่อตอบรับกระแสรักสุขภาพ หรือเลือกแนวทางปรับสูตรเครื่องดื่มที่มีอยู่เดิมด้วยการลดปริมาณน้ำตาล เช่น กลุ่มเครื่องดื่มชาที่ใช้สารทดแทนความหวานแทนน้ำตาล กลุ่มเครื่องดื่มน้ำผลไม้สูตรน้ำตาลน้อยและไม่ผสมน้ำตาล กลุ่มเครื่องดื่มน้ำอัดลมสูตรไม่มีน้ำตาล โดยการออกแบบบรรจุภัณฑ์ใหม่ให้เห็นความแตกต่างอย่างชัดเจนกับสูตรดั้งเดิม



ภาพที่ 1.11 ตัวอย่างผลิตภัณฑ์ที่ปรับสูตรลดน้ำตาล

1.3 ประเภทและระดับของผลิตภัณฑ์

ผลิตภัณฑ์ หมายถึง สิ่งใดก็ตามที่สามารถตอบสนองความต้องการและความจำเป็นของผู้บริโภคได้ เป็นส่วนผสมของลักษณะที่จับต้องได้ (Tangible) และจับต้องไม่ได้ (Intangible) ซึ่งผู้ซื้อจะยอมรับผลิตภัณฑ์นั้นถ้าสามารถตอบสนองความต้องการของผู้ซื้อได้ การแบ่งประเภทของผลิตภัณฑ์สามารถแบ่งได้หลายเกณฑ์ ดังนี้

1.3.1 การแบ่งประเภทผลิตภัณฑ์ตามความคงทน สามารถแบ่งได้เป็น 3 ประเภท ได้แก่

- 1) ผลิตภัณฑ์คงทน (Durable goods) เป็นผลิตภัณฑ์ที่มีอายุการใช้งานนาน ผู้บริโภคจึงมักต้องการการบริการเพิ่มเติม เช่น การรับประกัน ตัวอย่างผลิตภัณฑ์เช่น โต๊ะ เก้าอี้ เครื่องปรับอากาศ
- 2) ผลิตภัณฑ์ไม่คงทน (Nondurable goods) เป็นผลิตภัณฑ์ที่มีอายุการใช้งานสั้น ต้องซื้อบ่อย ตัวอย่างผลิตภัณฑ์ เช่น อาหาร เครื่องดื่ม
- 3) ผลิตภัณฑ์ในรูปของบริการ (Service) เป็นผลิตภัณฑ์ที่มีการผลิต การขาย และมีการบริโภคเกิดขึ้นภายในเวลาเดียวกัน ไม่สามารถแยกออกจากกันได้ (Inseparable) มีความหลากหลาย (Variable) และไม่สามารถเก็บไว้ได้ (Perishable) เช่น การตัดผม มีการให้บริการและรับบริการในเวลาเดียวกัน มีช่างตัดหลายคน มีทรงผมให้เลือกเยอะ และไม่สามารถจะเก็บความรู้สึกในขณะที่ได้รับบริการไว้ได้

1.3.2 การแบ่งประเภทผลิตภัณฑ์ตามลักษณะกายภาพ สามารถแบ่งได้เป็น 2 ประเภท ได้แก่

- 1) ผลิตภัณฑ์ที่จับต้องได้ (Tangible goods) อาจจะเป็นสินค้าที่คงทนหรือไม่คงทนก็ได้
- 2) ผลิตภัณฑ์ที่จับต้องไม่ได้ (Intangible goods) ผลิตภัณฑ์ประเภทนี้ต้องการการควบคุมคุณภาพเป็นพิเศษ และต้องสร้างความเชื่อถือต่อกันระหว่างผู้ขายและผู้ซื้อ

1.3.3 การแบ่งประเภทผลิตภัณฑ์ตามเกณฑ์ผู้ใช้ สามารถแบ่งได้เป็น 3 ประเภท ได้แก่

- 1) ผลิตภัณฑ์ทางการเกษตรและวัตถุดิบ (Agricultural product and raw material) เป็นสินค้าที่มีการเจริญเติบโต เพาะปลูกหรือได้มาจากพื้นดินหรือทะเล เช่น ธัญชาติ ธัญพืช ผัก ผลไม้ เหลือกและทราย โดยทั่วไปผลิตภัณฑ์ประเภทนี้จะขายเป็นปริมาณมากและมีมูลค่าต่อหน่วยน้อย
- 2) ผลิตภัณฑ์เพื่อผู้บริโภค (Consumer goods) คือ สินค้าหรือบริการที่ผู้ซื้อซื้อไปเพื่อการอุปโภคบริโภคเองในครอบครัวเพื่อตอบสนองความต้องการหรือสร้างความพึงพอใจแก่ตนเองหรือสมาชิกภายในครอบครัว แบ่งออกได้เป็น 4 ประเภทตามความพยายามในการตัดสินใจซื้อ ความถี่ในการซื้อ และคุณสมบัติของผลิตภัณฑ์ ได้แก่

- ผลิตภัณฑ์สะดวกซื้อ (Convenience products) เป็นผลิตภัณฑ์เพื่ออุปโภคบริโภคที่ผู้บริโภคซื้อบ่อยๆ ซื้อทันทีโดยไม่ต้องวางแผนล่วงหน้า ไม่ต้องเปรียบเทียบผลิตภัณฑ์ ไม่เสียเวลาในการไตร่ตรองนาน เป็นผลิตภัณฑ์ที่มักใช้ในชีวิตรประจำวัน ความภักดีต่อตราสินค้าค่อนข้างต่ำ ดังนั้นหากไม่มีสินค้าจำหน่ายผู้บริโภคจะหันไปซื้อผลิตภัณฑ์ยี่ห้ออื่นแทน ผลิตภัณฑ์ประเภทนี้หาซื้อได้ง่าย เช่น สบู่ ยาสีฟัน

- ผลิตภัณฑ์เลือกซื้อ (Shopping products) หรือ ผลิตภัณฑ์เปรียบเทียบซื้อ ผู้บริโภคมีความถี่ในการซื้อผลิตภัณฑ์ประเภทนี้ต่ำกว่าผลิตภัณฑ์สะดวกซื้อ ในการซื้อผู้บริโภคจะทำการหาข้อมูลเพื่อทำการเปรียบเทียบผลิตภัณฑ์ทั้งด้านราคา คุณภาพ รูปแบบ การบริการหลังการขาย เช่น เครื่องใช้ไฟฟ้า เสื้อผ้า

- ผลิตภัณฑ์เจาะจงซื้อ (Specialty products) เป็นผลิตภัณฑ์เพื่ออุปโภคบริโภคซึ่งมีคุณลักษณะพิเศษ มีลักษณะโดดเด่นเป็นเอกลักษณ์ ทำให้มีกลุ่มผู้ซื้อเฉพาะ มีความพยายามในการซื้อเป็นพิเศษ เช่น รถยนต์ เครื่องเสียง นาฬิกา

- ผลิตภัณฑ์ที่ไม่ได้แสวงซื้อ (Unsought products) เป็นผลิตภัณฑ์ที่ผู้บริโภคไม่มีความรู้เกี่ยวกับผลิตภัณฑ์นี้มาก่อน ไม่สนใจจะซื้อจนกว่าจะมีคนมาแนะนำเสนอขาย เช่น การทำประกันชีวิต

3) ผลิตภัณฑ์เพื่ออุตสาหกรรม (Industrial goods) คือ สินค้าหรือบริการที่ผู้ซื้อซื้อไปใช้งานภายในองค์กร หรือซื้อไปใช้ในกระบวนการผลิตสินค้าหรือบริการแล้วนำออกมาขายเพื่อแสวงหากำไร แบ่งออกได้เป็น 4 ประเภท ได้แก่

- วัตถุดิบ (Materials) หมายถึง วัตถุดิบและชิ้นส่วนประกอบในกระบวนการผลิต
- สินค้าคงทน (Capitalism) เป็นสินค้าคงทนอยู่ในส่วนของการผลิต เช่น ตัวอาคาร โรงงาน อุปกรณ์เครื่องจักร
- อะไหล่และบริการเสริม (Supplied and services) เป็นวัสดุสำนักงานที่มีอายุการใช้งานสั้น และการบริการเสริมเพื่อให้การผลิตดำเนินต่อไปได้ เช่น การดูแลรักษา ซ่อมแซมอุปกรณ์
- ธุรกิจบริการ (Business services) เป็นบริการที่ช่วยลูกค้าในการดำเนินกิจการ เช่น บริการที่ปรึกษาทางด้านกฎหมาย ด้านภาษี

1.3.4 ระดับของผลิตภัณฑ์ 5 ระดับ (Five levels of product) หรือ องค์ประกอบของผลิตภัณฑ์ (Product component) เป็นการพิจารณาถึงคุณสมบัติหรือคุณค่าของผลิตภัณฑ์ที่นักพัฒนาผลิตภัณฑ์ต้องนำมาใช้พิจารณากำหนดลักษณะผลิตภัณฑ์ที่สามารถตอบสนองความต้องการของผู้บริโภคได้ สามารถแบ่งคุณค่าของผลิตภัณฑ์ที่มีต่อผู้บริโภคออกเป็น 5 ระดับ ได้แก่

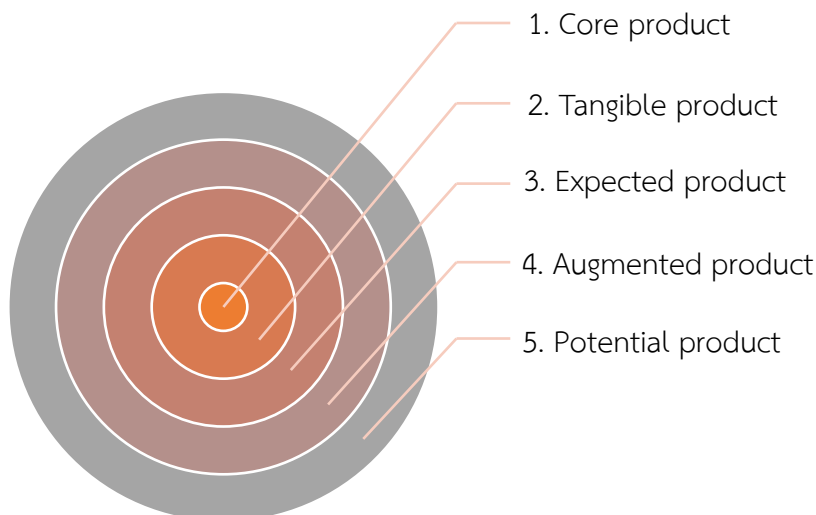
ระดับที่ 1 ผลิตภัณฑ์หลัก (Core product) หมายถึง ประโยชน์พื้นฐานของผลิตภัณฑ์ที่ผู้บริโภคได้รับจากการซื้อสินค้าโดยตรง เช่น ความสามารถในการติดของกาว

ระดับที่ 2 ผลิตภัณฑ์ที่จับต้องได้ (Tangible product) หมายถึง ลักษณะทางกายภาพที่ผู้บริโภคสามารถสัมผัสหรือรับรู้ได้ ซึ่งเป็นส่วนที่เสริมผลิตภัณฑ์ให้ทำหน้าที่สมบูรณ์ขึ้นหรือเชิญชวนให้ใช้ยิ่งขึ้น ประกอบไปด้วยคุณภาพ รูปแบบ บรรจุภัณฑ์ และตราสินค้า เช่น กาวแท่งยี่ห้อ 3M

ระดับที่ 3 ผลิตภัณฑ์คาดหวัง (Expected product) หมายถึง ประโยชน์หรือเงื่อนไขที่ผู้บริโภคคาดว่าจะได้รับจากผลิตภัณฑ์ เช่น กาวแท่งใช้ง่าย สะดวก ติดติดสามารถลอกออกมาติดใหม่ได้ ไม่เลอะมือ

ระดับที่ 4 ผลิตภัณฑ์ควบ (Augmented product) หมายถึง ประโยชน์เพิ่มเติม หรือบริการที่ผู้ซื้อจะได้รับควบคู่กับการซื้อสินค้า เช่น บริการก่อนและหลังการขาย การติดตั้ง การขนส่ง การรับประกัน การให้สินเชื่อ และการให้บริการ เช่น ซื้อกาวแท่ง 3 อันแถมฟรี 1 อัน

ระดับที่ 5 ผลิตภัณฑ์ที่มีศักยภาพ (Potential product) หมายถึง คุณสมบัติของผลิตภัณฑ์ใหม่ที่มีการเปลี่ยนแปลงหรือพัฒนาไปเพื่อสนองความต้องการของลูกค้าในอนาคต เช่น บรรจุภัณฑ์ของกาวแท่งสามารถย่อยสลายได้เองในธรรมชาติ หรือสามารถลอกกลากบรรจุภัณฑ์เพื่อส่งชิงโชค



ภาพที่ 1.12 ระดับของผลิตภัณฑ์ 5 ระดับ

1.4 ความหมายและประเภทของผลิตภัณฑ์ใหม่

การพัฒนาผลิตภัณฑ์ใหม่ เป็นการพัฒนาและแนะนำผลิตภัณฑ์ซึ่งบริษัทไม่เคยมีการผลิตมาก่อนเข้าสู่ตลาด หรือเป็นการเสนอผลิตภัณฑ์เดิมสู่ตลาดใหม่ซึ่งบริษัทไม่เคยเข้าสู่ตลาดกลุ่มนี้มาก่อน

1.4.1 ผลิตภัณฑ์ใหม่ (New product) หมายถึง ผลิตภัณฑ์ที่มีลักษณะ ดังนี้

- (1) มีลักษณะริเริ่ม หรือมีความเป็นผลิตภัณฑ์นวัตกรรม (Innovative products) กล่าวคือ เป็นผลิตภัณฑ์ที่มีลักษณะแตกต่างจากผลิตภัณฑ์เดิมที่มีอยู่ โดยสามารถสนองความต้องการเดิมหรือความต้องการใหม่ของผู้ซื้อได้ เช่น มือถือที่สามารถใช้งานได้อย่างทั้งพุดคุย ถ่ายรูป เล่นอินเทอร์เน็ต
- (2) ผลิตภัณฑ์ใหม่ที่สามารถทดแทนผลิตภัณฑ์เดิม แต่มีลักษณะแตกต่างจากผลิตภัณฑ์เดิมอย่างเห็นได้ชัด หรือมีการปรับปรุงใหม่ (Modified products) เช่น ผลิตภัณฑ์น้ำยาซักผ้าใช้แทนผงซักฟอก
- (3) ผลิตภัณฑ์เลียนแบบ (Imitative or Me-too products) เป็นผลิตภัณฑ์ที่ใหม่สำหรับบริษัทแต่เก่าสำหรับตลาดสินค้านั้น เป็นการเลียนแบบ เช่น ผลิตภัณฑ์เค้กโตเกียวบานาน่าซึ่งเป็นของฝากที่นิยมจากประเทศญี่ปุ่น และมีผู้ประกอบการคนไทยนำมาทำเลียนแบบและจำหน่ายในไทย

1.4.2 การแบ่งประเภทของผลิตภัณฑ์ใหม่ แบ่งได้เป็น 7 ประเภท ดังนี้

- (1) Line extension เป็นผลิตภัณฑ์ที่เกิดจากการขยายสายการผลิต ไม่จำเป็นต้องอาศัยเทคโนโลยีใหม่ๆ นโยบายทางการตลาดเหมือนเดิม เช่น การเปลี่ยนแปลงรสชาติและสารให้กลิ่นในนมพร้อมดื่ม ผลิตภัณฑ์แครกเกอร์สูตรเดิมแต่ลดขนาดตัวผลิตภัณฑ์ลง



ภาพที่ 1.13 ตัวอย่างผลิตภัณฑ์ประเภท Line extension

(2) Repositioned existing products เป็นผลิตภัณฑ์ที่เกิดจากการเปลี่ยนตำแหน่งทางการตลาดของผลิตภัณฑ์ มีวัตถุประสงค์ในการใช้งานที่แตกต่างจากที่บริษัทได้กำหนดไว้ เช่น วิตามินบางกลุ่มนำมายาอูทัยทิพย์ซึ่งมีสีแดงและปกติบริโภคโดยใส่น้ำดื่มแต่กลับนำมาทาแก้มและปากให้มีสีแดง ผลิตภัณฑ์โยเกิร์ตนำมาใช้พอกหน้าแทนการรับประทาน



ภาพที่ 1.14 ตัวอย่างผลิตภัณฑ์ประเภท Repositioned existing products

(3) New form of existing products เป็นผลิตภัณฑ์รูปแบบใหม่ เช่น ผงปรุงรสคนอร์จากเดิมเป็นก้อนอาจเปลี่ยนเป็นผง ผงซักฟอกจากในรูปผงเปลี่ยนมาเป็นรูปก้อนหรือของเหลว และยากันยุงรูปแบบต่างๆ



ภาพที่ 1.15 ตัวอย่างผลิตภัณฑ์ประเภท New form of existing products

(4) Reformulation of existing products เป็นผลิตภัณฑ์ที่เกิดจากการเปลี่ยนสูตร เพื่อลดต้นทุนการผลิต เพื่อเพิ่มคุณค่าทางโภชนาการ เป็นการหาวัตถุดิบใหม่มาทดแทนวัตถุดิบเดิม เช่น ขนมขบเคี้ยวจากข้าวไรซ์เบอร์รี่ ขนมขบเคี้ยวเสริมโปรตีน น้ำสลัดสูตรไม่มีคอเลสเตอรอล



ภาพที่ 1.16 ตัวอย่างผลิตภัณฑ์ประเภท Reformulation of existing products

(5) New packaging of existing products เป็นผลิตภัณฑ์ที่เกิดจากการเปลี่ยนภาชนะบรรจุเพื่อดึงดูดความสนใจ เพื่อให้ง่ายต่อการใช้งานและบริโภค เช่น สบู่ แชมพู น้ำอัดลม น้ำจิ้มไก่ น้ำพริกเผา



ภาพที่ 1.17 ตัวอย่างผลิตภัณฑ์ประเภท New packaging of existing products

(6) Innovative and value added products เป็นผลิตภัณฑ์ที่เกิดจากการเปลี่ยนแปลงในตัวผลิตภัณฑ์เพื่อเพิ่มมูลค่า เช่น ไข่ลวกพร้อมรับประทาน น้ำแกงพร้อมปรุงบรรจุกล่อง คางกุ้งทอดปรุงรส ทุ่นำกระป๋องคู่กับแครกเกอร์ กาแฟผสมน้ำตาลและครีมเทียมในซองเดียวกัน (3 in 1) วุ้นเส้นขายพร้อมน้ำปรุงรส สำหรับทำวุ้นเส้นอบหม้อดิน แผลงอบกรอบ



ภาพที่ 1.18 ตัวอย่างผลิตภัณฑ์ประเภท Innovative and value added products

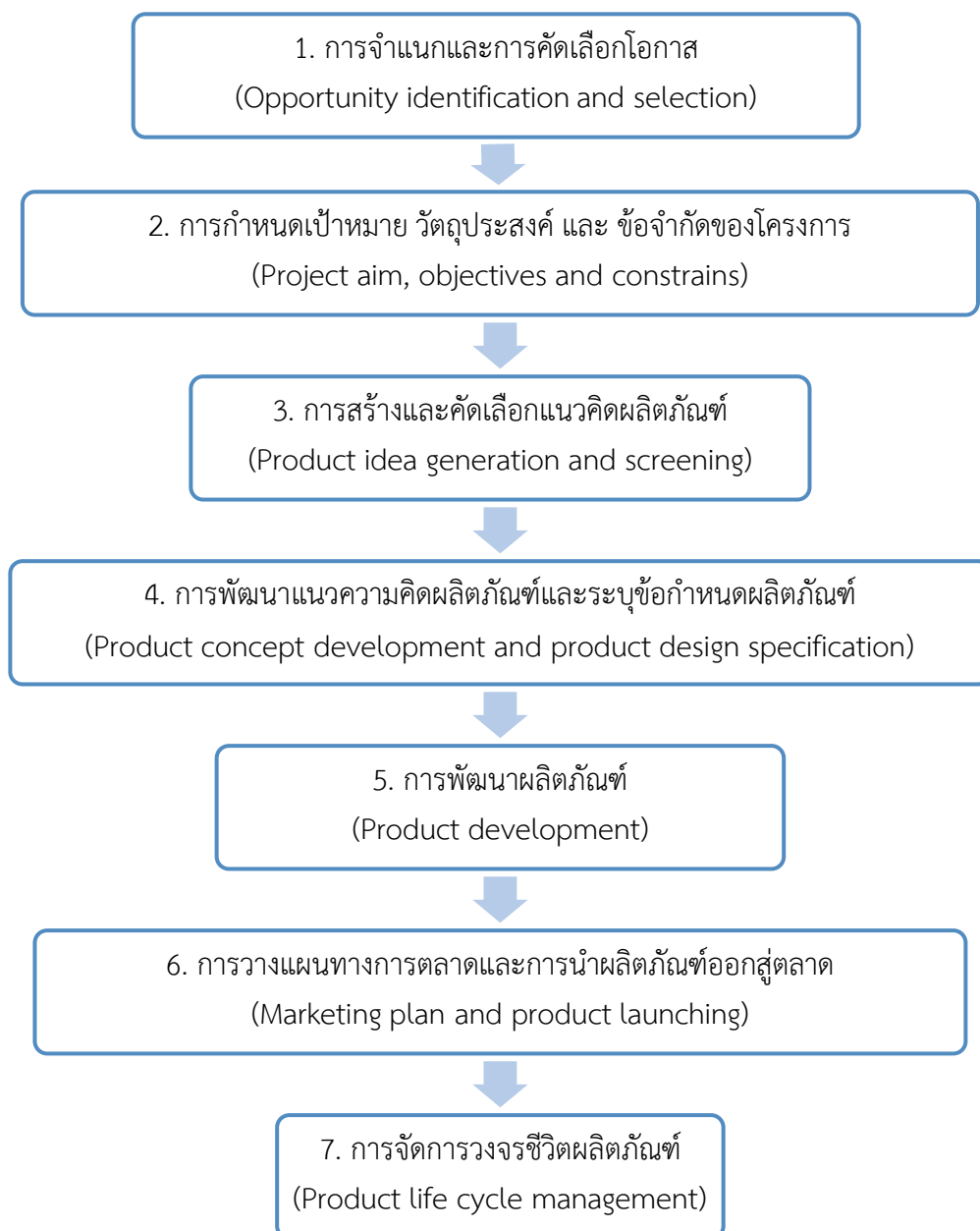
(7) Creative products เป็นผลิตภัณฑ์ที่มีการคิดค้นขึ้นใหม่ไม่เคยมีมาก่อนในท้องตลาด เกิดจากความคิดสร้างสรรค์ และการพัฒนาทางด้านเทคโนโลยี เช่น เบอร์เกอร์หมูจากพืช หมูกรอบจากพืช ช็อกโกแลตแท่งวีแกน โยเกิร์ตเยลลี่ ไอศกรีมมะม่วงน้ำปลาหวาน ผลิตภัณฑ์มะพร้าวขึ้นหวาน และ ส้มตำไทย อบกรอบเพียงฉีกซองเติมน้ำสามารถรับประทานได้ทันที



ภาพที่ 1.19 ตัวอย่างผลิตภัณฑ์ประเภท Creative products

1.5 กระบวนการพัฒนาผลิตภัณฑ์

กระบวนการพัฒนาผลิตภัณฑ์ แบ่งออกได้เป็น 7 ขั้นตอน ได้แก่ 1) การจำแนกและคัดเลือกโอกาส 2) การกำหนดเป้าหมาย วัตถุประสงค์ และข้อจำกัดของโครงการ 3) การสร้างและคัดเลือกแนวคิดผลิตภัณฑ์ 4) การพัฒนาแนวความคิดผลิตภัณฑ์และการระบุข้อกำหนดผลิตภัณฑ์ 5) การพัฒนาผลิตภัณฑ์ 6) การวางแผนทางการตลาดและการนำผลิตภัณฑ์ออกสู่ตลาด และ 7) การจัดการวงจรชีวิตผลิตภัณฑ์ ดังแสดงในภาพที่ 1.20



ภาพที่ 1.20 กระบวนการพัฒนาผลิตภัณฑ์

ขั้นตอนกระบวนการพัฒนาผลิตภัณฑ์ อธิบายได้โดยสังเขป ดังนี้

ขั้นที่ 1 การจำแนกและคัดเลือกโอกาส

ก่อนที่จะเริ่มโครงการพัฒนาผลิตภัณฑ์ นักวิจัยต้องวิเคราะห์โอกาสหรือความเป็นไปได้ที่ผลิตภัณฑ์จะประสบความสำเร็จ ต้องสามารถผลิตได้จริงและขายได้กำไร เป็นที่ยอมรับของผู้บริโภค โดยการรวบรวมข้อมูลปฐมภูมิและทุติยภูมิใน 3 ด้านได้แก่ การผลิต การตลาด และ ผู้บริโภค โดยปกติเน้นแรงผลักดันที่ส่งผลให้ต้องมีการพัฒนาผลิตภัณฑ์เกิดจากการเปลี่ยนแปลง 6 ด้าน ได้แก่ ด้านธุรกิจ ด้านเศรษฐกิจ ด้านสังคม ด้านการผลิต ด้านเทคโนโลยี และด้านการเปลี่ยนแปลงที่มีผลต่อสิ่งแวดล้อม นักพัฒนาผลิตภัณฑ์ต้องทำการรวบรวมข้อมูลดังกล่าวเพื่อนำมาใช้วิเคราะห์ค้นหาโอกาสของผลิตภัณฑ์ใหม่ว่ามีแนวโน้มสามารถผลิตได้และขายได้จริงหรือไม่

ในการวิเคราะห์โอกาสของผลิตภัณฑ์ใหม่ สามารถแบ่งพิจารณาออกได้เป็น 3 ด้าน ได้แก่

1) การวิเคราะห์สถานการณ์ (Situation analysis) ทำได้โดยวิธี SWOT analysis คือ การวิเคราะห์สภาพองค์กรหรือหน่วยงานในปัจจุบัน โดยการพิจารณาสภาพการณ์ภายในและภายนอก เพื่อค้นหาจุดแข็ง (Strengths, S) และจุดอ่อน (Weaknesses, W) จากปัจจัยภายในเพื่อให้รู้ตัวเรา และค้นหาโอกาส (Opportunities, O) และอุปสรรค (Threats, T) จากปัจจัยภายนอกเพื่อให้รู้เขา ดังภาษิตจีนโบราณกล่าวว่า “รู้เขารู้เรา รบร้อยครั้งชนะร้อยครั้ง”



ภาพที่ 1.21 ตัวอย่างการวิเคราะห์ SWOT analysis ของร้านอาหารที่ต้องการย้ายทำเลที่ตั้ง

2) แรงกระตุ้นจากกลยุทธ์ของบริษัท (Input from corporate strategy) ได้แก่ (1) เทคโนโลยีที่บริษัทมี เช่น บริษัทมีเครื่องบรรจุกระป๋องที่มีประสิทธิภาพในการผลิตสูง, (2) ส่วนแบ่งตลาดที่บริษัทถือครอง เช่น เน้นผลิตภัณฑ์อาหารเพื่อสุขภาพ และ (3) แนวปฏิบัติของบริษัท เช่น ผลิตเฉพาะผลิตภัณฑ์เพื่อสุขภาพ

3) แรงกระตุ้นจากแหล่งอื่น (Input from other source) ได้แก่ (1) ทรัพยากรที่ยังไม่ได้มีการนำมาใช้ โดยพิจารณาทั้งในด้านเทคโนโลยี การเงิน ผลิตภัณฑ์ และตลาด, (2) ทรัพยากรใหม่ ซึ่งได้จากการค้นพบ การเข้าซื้อกิจการ การขยายตลาด และ (3) แรงผลักดันภายนอก เช่น คุณภาพของผลิตภัณฑ์คู่แข่ง ข้อกำหนดกฎหมายมาตรฐานที่ต้องปฏิบัติ

ขั้นที่ 2 การกำหนดเป้าหมาย วัตถุประสงค์ และข้อจำกัดของโครงการ

ก่อนเริ่มโครงการพัฒนาผลิตภัณฑ์นั้น ทีมฝ่ายวิจัยและพัฒนาต้องระบุ ภาพรวมการทำงานทั้งหมดให้ชัดเจน ในแต่ละขั้นตอนของการดำเนินงานผลลัพธ์ที่เกิดขึ้นคืออะไร และสิ่งใดเป็นข้อจำกัดของโครงการ ซึ่งทุกคนที่มีส่วนเกี่ยวข้องกับการพัฒนาผลิตภัณฑ์ของโครงการจะต้องร่วมกันอภิปรายและกำหนดเป้าหมาย (Aim) วัตถุประสงค์ (Objectives) และข้อจำกัดของโครงการ (Constraints)

1) เป้าหมายของโครงการ เป็นผลลัพธ์สุดท้ายที่จะได้รับเมื่อเสร็จสิ้นโครงการ การกำหนดเป้าหมายต้องมีความชัดเจน ไม่กำกวม ไม่ซับซ้อน และต้องมีความจำเพาะเจาะจง สอดคล้องกับกลยุทธ์ของบริษัท โครงการพัฒนาผลิตภัณฑ์ทุกโครงการต้องเริ่มต้นด้วยการกำหนดเป้าหมาย ซึ่งมีการเขียนเป็นลายลักษณ์อักษร โดยมีการเห็นพ้องต้องกันของผู้ที่เกี่ยวข้องในโครงการ

2) วัตถุประสงค์ เป็นจุดหมายปลายทางของขั้นตอนต่างๆ ของโครงการ สร้างได้จากผลลัพธ์หลักที่จะเกิดขึ้นในแต่ละขั้นตอน ซึ่งไม่ควรมีหลายข้อเกินไป วัตถุประสงค์ที่มีผลต่อความสำเร็จของโครงการควรระบุให้ละเอียดเพื่อใช้เป็นแนวทางในการกำหนดขั้นตอนการทำงาน ทั้งนี้วัตถุประสงค์ที่กำหนดต้องวัดค่าได้ เพื่อจะได้ประเมินว่าโครงการบรรลุวัตถุประสงค์ที่กำหนดไว้หรือไม่

3) การระบุข้อจำกัดของโครงการ คือ การกำหนดปัจจัย หรือ จำกัดขอบเขตของโครงการ ซึ่งต้องมีการกำหนดให้ชัดเจนก่อนเริ่มทำการพัฒนาผลิตภัณฑ์ ข้อจำกัดที่ต้องระบุได้แก่ ข้อจำกัดทางด้านผลิตภัณฑ์ ด้านกระบวนการแปรรูป ด้านการตลาด ด้านการเงิน ด้านบริษัท และด้านสิ่งแวดล้อม

ภาพที่ 1.22 แสดงตัวอย่างการกำหนดเป้าหมาย วัตถุประสงค์ และข้อจำกัดของโครงการ ของโครงการพัฒนาผลิตภัณฑ์เค้กกล้วยหอมไขมันต่ำ



ภาพที่ 1.22 ตัวอย่างการกำหนดเป้าหมาย วัตถุประสงค์ และข้อจำกัดของโครงการ

ขั้นที่ 3 การสร้างและคัดเลือกแนวคิดผลิตภัณฑ์

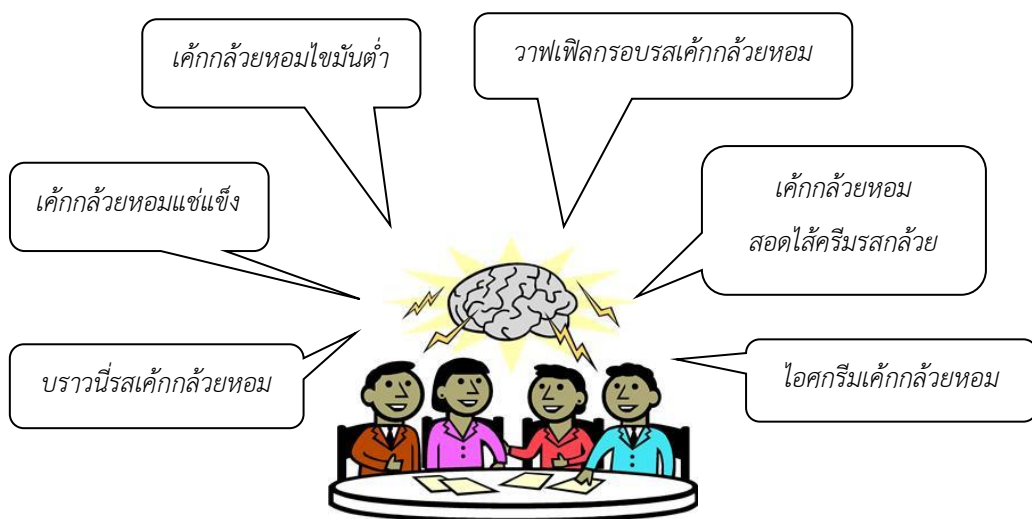
(1) การสร้างแนวคิดผลิตภัณฑ์ (Product idea generation) เป็นกระบวนการสร้างสรรค์ความคิดอย่างเป็นระบบ โดยการนำความรู้ด้านผู้บริโภค การตลาด เทคโนโลยี และสิ่งแวดล้อม มาประยุกต์ใช้ในการสร้างผลิตภัณฑ์ใหม่ ในขั้นตอนการสร้างแนวคิดผลิตภัณฑ์ใหม่จำเป็นต้องทำการวิจัยเพื่อค้นหาความคิดผลิตภัณฑ์ที่ผู้บริโภคต้องการและมีความเป็นไปได้ทางการตลาด ตลอดจนพยายามค้นหาการเปลี่ยนแปลงหรือสร้างสิ่งใหม่ให้แก่ผลิตภัณฑ์ โดยความคิดใหม่นี้จะต้องสอดคล้องกับกลยุทธ์ของบริษัทและสามารถพัฒนาผลิตภัณฑ์ได้จริงตาม วัตถุประสงค์ เป้าหมาย และข้อจำกัดของโครงการ

การค้นหาข้อมูลเพื่อสร้างแนวคิดผลิตภัณฑ์ใหม่สามารถสืบค้นได้จากข้อมูลทางตรงหรือฐานข้อมูลที่มีอยู่แล้ว เช่น ข้อมูลของฝ่ายการตลาดและฝ่ายขาย เอกสารสิทธิบัตร ฐานข้อมูลออนไลน์ นอกจากนี้ยังสามารถค้นหาและวิจัยเพื่อรวบรวมข้อมูลมาสร้างแนวคิดผลิตภัณฑ์ใหม่ด้วยตนเอง โดยอาศัยเทคนิคต่างๆ เช่น การสัมภาษณ์ การสำรวจ และการจดบันทึก ภาพที่ 1.23 แสดงปัจจัยที่ส่งผลต่อการสร้างแนวคิดผลิตภัณฑ์



ภาพที่ 1.23 ปัจจัยที่ควรพิจารณาในการสร้างแนวคิดผลิตภัณฑ์ใหม่

วิธีการสร้างแนวคิดผลิตภัณฑ์ แบ่งออกได้เป็น 2 วิธีใหญ่ ได้แก่ 1) วิธีการสร้างแนวคิดด้วยตนเอง เช่น เทคนิคการคิดแนวข้าง (Lateral thinking) และ 2) วิธีการสร้างความคิดผลิตภัณฑ์แบบกลุ่ม เช่น การอภิปรายกลุ่ม (Focus group discussion), การวิเคราะห์ช่องว่างในตลาด (Gap analysis), การวิเคราะห์ลักษณะของผลิตภัณฑ์ (Attribute analysis), การวิเคราะห์ความต้องการของผู้บริโภค (Consumer needs analysis), การวิเคราะห์มิติของผลิตภัณฑ์ (Grid analysis), การคิดสร้างสรรค์อย่างเป็นระบบ (Systematic inventive thinking) และการสัมภาษณ์เชิงลึกรายบุคคล (Individual depth interview) สิ่งสำคัญที่นักพัฒนาผลิตภัณฑ์ต้องคำนึงถึงในการสร้างแนวคิดผลิตภัณฑ์ คือ ผู้บริโภคเป้าหมายเป็นใคร และปัญหาหลักที่ผู้บริโภคพบเจอในผลิตภัณฑ์และต้องการแก้ไข (Consumer pain points) คืออะไร



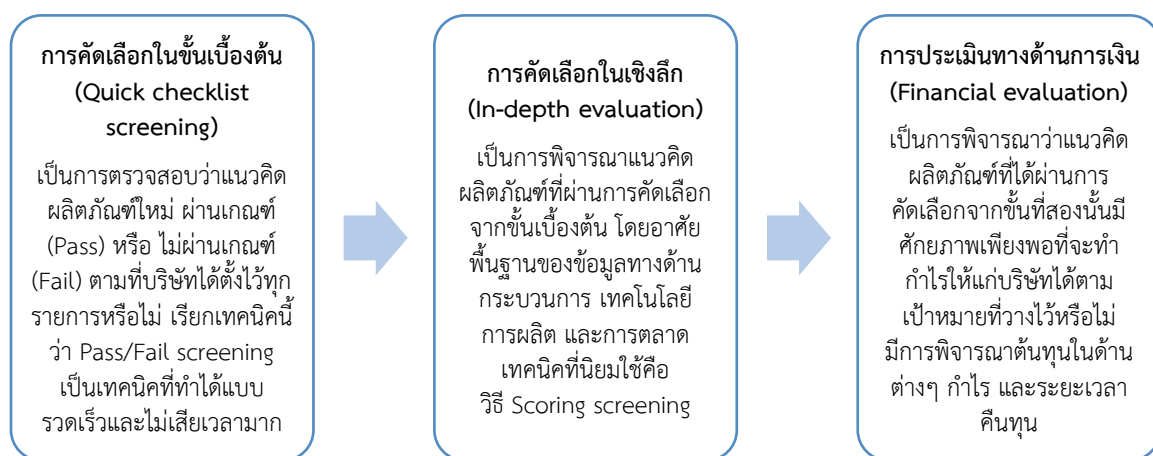
ภาพที่ 1.24 ตัวอย่างแนวคิดผลิตภัณฑ์เพื่อพัฒนาผลิตภัณฑ์เค็กกล้วยหอม

วิธีทางสถิติที่นิยมใช้ในขั้นตอนการสร้างแนวคิดผลิตภัณฑ์ เช่น การวิเคราะห์เชิงพรรณนา (Descriptive analysis), การวิเคราะห์ตัวแปรหลายตัว (Multivariate analysis) เช่น การวิเคราะห์แบบหลายมิติ (Multidimensional scaling method), การวิเคราะห์ปัจจัย (Factor analysis) และ การวิเคราะห์กลุ่ม (Cluster analysis)

(2) การคัดเลือกแนวคิดผลิตภัณฑ์ (Product idea screening) คือ การได้มาซึ่งแนวคิดผลิตภัณฑ์ที่มีโอกาสประสบความสำเร็จ และตัดแนวคิดผลิตภัณฑ์ที่อาจจะไม่ประสบความสำเร็จออกไปโดยใช้ข้อมูลด้านต่างๆ มาพิจารณา การคัดเลือกแนวคิดผลิตภัณฑ์สามารถทำได้โดยอาศัยพื้นฐานความรู้และประสบการณ์ของที่ปรึกษาบริษัท บุคคลภายในบริษัท ร่วมกับ ข้อมูลจากแหล่งภายในและนอกบริษัท

ปัจจัยที่ใช้ในการคัดเลือกแนวคิดผลิตภัณฑ์ควรพิจารณาจากกลยุทธ์ของบริษัท เป้าหมาย และข้อจำกัดของโครงการ ทั้งนี้การคัดเลือกปัจจัยเพื่อใช้เป็นเกณฑ์ในการคัดเลือกแนวคิดผลิตภัณฑ์ของแต่ละบริษัทและแต่ละผลิตภัณฑ์มีความแตกต่างกัน เช่น เงินทุน คือ ปัจจัยอันดับแรกที่บริษัทขนาดเล็กมักจะคำนึงถึงและใช้ในการคัดเลือกแนวคิดผลิตภัณฑ์ใหม่ การเลือกปัจจัยเพื่อใช้ในการคัดเลือกแนวคิดผลิตภัณฑ์ขึ้นอยู่กับ นโยบายของบริษัท ทรัพยากรของบริษัท สภาพแวดล้อมของบริษัท และระดับความเป็นนวัตกรรมของผลิตภัณฑ์ เนื่องจากปัจจัยที่ใช้ในการคัดเลือกแนวคิดผลิตภัณฑ์มีเป็นจำนวนมาก จึงจำเป็นต้องเลือกโดยการเรียงลำดับความสำคัญของปัจจัย (สำคัญมากที่สุด, สำคัญมาก, สำคัญปานกลาง, สำคัญน้อย และไม่สำคัญ) ปัจจัยที่จำเป็นและมีความเสี่ยงต่อความสำเร็จของผลิตภัณฑ์จะได้รับการคัดเลือกในอันดับต้นๆ

ขั้นตอนการคัดเลือกแนวคิดผลิตภัณฑ์ แบ่งออกได้เป็น 3 ขั้นตอน คือ 1) การคัดเลือกในขั้นเบื้องต้น (Quick checklist screening) 2) การคัดเลือกในเชิงลึก (In-depth evaluation) และ 3) การประเมินทางด้านการเงิน (Financial evaluation) ดังภาพที่ 1.25



ภาพที่ 1.25 ขั้นตอนการคัดเลือกแนวคิดผลิตภัณฑ์

ขั้นที่ 4 การพัฒนาแนวความคิดผลิตภัณฑ์ และการระบุข้อกำหนดผลิตภัณฑ์

(1) การพัฒนาและทดสอบแนวความคิดผลิตภัณฑ์

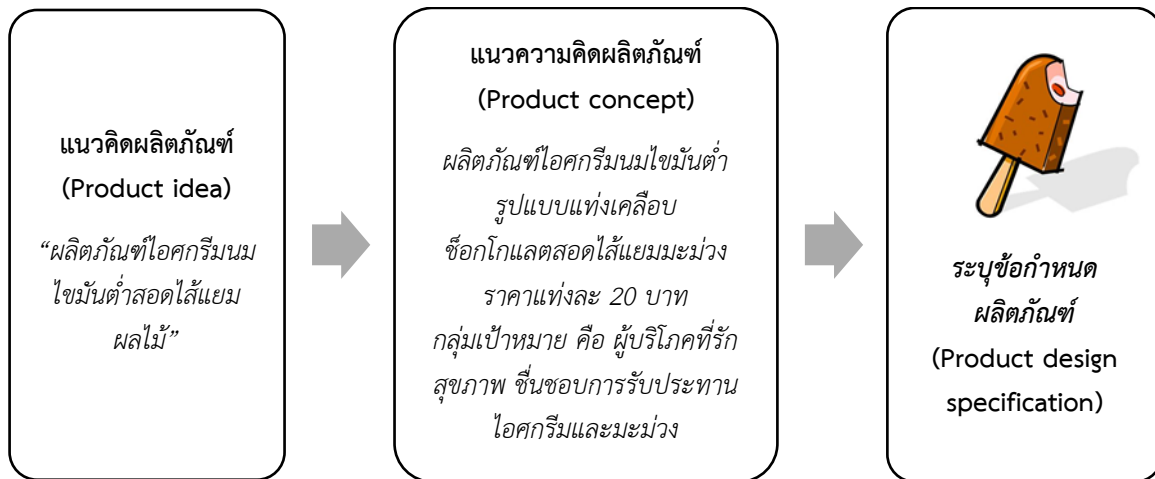
การพัฒนาและทดสอบแนวความคิดผลิตภัณฑ์ เป็นการนำแนวคิดผลิตภัณฑ์ (Product idea) ที่ได้ผ่านการคัดเลือกแล้วมาพัฒนาเป็นแนวความคิดผลิตภัณฑ์ (Product concept) โดยระบุลักษณะของผลิตภัณฑ์ ประโยชน์ที่ผู้บริโภคต้องการจากผลิตภัณฑ์ 4 ด้านได้แก่ ด้านตัวผลิตภัณฑ์ (Product benefits), ด้านบรรจุภัณฑ์ (Package benefits), ด้านการใช้งาน (Use benefits) และด้านจิตใจ (Psychological benefits) รวมทั้งตำแหน่งทางการตลาดของผลิตภัณฑ์ ต่อจากนั้นนำแนวความคิดผลิตภัณฑ์ที่ถูกพัฒนาไปทดสอบกับกลุ่มผู้บริโภคเป้าหมายเพื่อดูความเป็นไปได้หรือแนวโน้มที่ผู้บริโภคสนใจ (Product concept validation) ซึ่งจะสามารถคาดคะเนปริมาณการผลิตและวางแผนทางการตลาดได้ ในการทดสอบอาจใช้วิธีการสำรวจความคิดเห็นของผู้บริโภคโดยใช้แบบสอบถามเป็นเครื่องมือ

การเขียนแนวความคิดผลิตภัณฑ์ที่ดีต้องสั้นได้ใจความ ใช้ข้อความ 3-4 ประโยค บรรยายให้เห็นภาพของผลิตภัณฑ์ โดยหลีกเลี่ยงการใช้ศัพท์ทางเทคนิค ใช้ภาษาที่เข้าใจง่าย ดึงดูดให้ผู้บริโภคสนใจ เริ่มต้นประโยคด้วยประเภทของผลิตภัณฑ์ ระบุชนิดของผลิตภัณฑ์ด้วยคำสั้นๆ คำแรกของการบรรยาย ให้รายละเอียดที่ถูกต้องของผลิตภัณฑ์ อย่างกล่าวอ้างคุณภาพของผลิตภัณฑ์เกินจริง แนวความคิดผลิตภัณฑ์จะต้องเชื่อถือได้และสมเหตุสมผล อธิบายวิธีการใช้และคุณค่าของผลิตภัณฑ์ แสดงให้ผู้บริโภคเห็นว่าผลิตภัณฑ์มีประโยชน์และมีข้อดีเหนือผลิตภัณฑ์อื่นที่มีอยู่ในท้องตลาด

ในขั้นตอนนี้อาจมีจำนวนแนวความคิดผลิตภัณฑ์มากกว่า 1 แนวความคิด จึงต้องทำการทดสอบและคัดเลือกแนวความคิดผลิตภัณฑ์ที่เหมาะสมที่สุด ผู้บริโภคจะเป็นผู้บอกว่าแนวความคิดใดดีและน่าสนใจมากที่สุด ต่อจากนั้นฝ่ายวิจัยของบริษัทจะทำหน้าที่พิจารณาว่าแนวความคิดนั้นสามารถผลิตได้และขายได้จริงหรือไม่โดยการพยากรณ์ยอดขาย กำไร และส่วนแบ่งทางการตลาด (Market share)

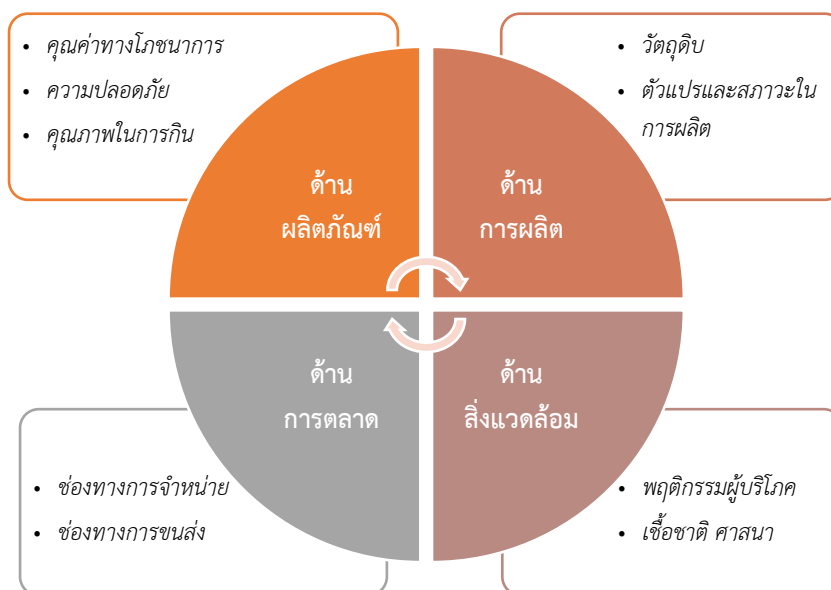
(2) การระบุข้อกำหนดผลิตภัณฑ์ (Product design specification)

ในอุตสาหกรรมอาหารนั้น การออกแบบข้อกำหนดผลิตภัณฑ์ คือ การกำหนดมาตรฐานคุณภาพผลิตภัณฑ์ (Product quality standard) ซึ่งกำหนดได้จากมาตรฐานตามกฎหมายกำหนดและจากขอบเขตค่าเป้าหมายที่ผู้บริโภคให้การยอมรับ ค่าพารามิเตอร์ที่ได้จากการออกแบบข้อกำหนดจะเป็นช่วงระหว่างค่าน้อยสุดและค่ามากที่สุดที่ยอมรับได้ การวัดค่าพารามิเตอร์เช่นนี้สามารถนำไปใช้ในการเปรียบเทียบผลิตภัณฑ์ใหม่กับผลิตภัณฑ์คู่แข่งได้เป็นอย่างดี ดังนั้นนักพัฒนาผลิตภัณฑ์ต้องมีความรู้ในศาสตร์ด้านการวิเคราะห์คุณภาพร่วมกับการวิเคราะห์ข้อมูลทางสถิติ



ภาพที่ 1.26 ตัวอย่างลำดับขั้นตอนการพัฒนาแนวความคิดผลิตภัณฑ์ และการระบุข้อกำหนดผลิตภัณฑ์

การระบุข้อกำหนดผลิตภัณฑ์ต้องเตรียมข้อมูล 4 ด้าน ได้แก่ ด้านตัวผลิตภัณฑ์ ด้านกระบวนการผลิต ด้านการตลาด และ ด้านสิ่งแวดล้อม ดังภาพที่ 1.27 โดยข้อกำหนดผลิตภัณฑ์ประกอบด้วย 3 ส่วน ได้แก่ 1) แนวความคิดผลิตภัณฑ์ (General statement of product concept), 2) คุณลักษณะของผลิตภัณฑ์ (Design characteristics) ที่ผู้บริโภคต้องการ เช่น ลักษณะทางด้านประสาทสัมผัส และ 3) ข้อกำหนดผลิตภัณฑ์ (Design specification) เช่น ข้อกำหนดด้านวัตถุดิบ, ข้อกำหนดด้านกระบวนการผลิต, ข้อกำหนดด้านบรรจุภัณฑ์, ข้อกำหนดด้านการตลาด, สภาพแวดล้อม และค่าใช้จ่ายที่ใช้ในกระบวนการผลิตและการตลาด



ภาพที่ 1.27 ตัวอย่างข้อมูลที่ต้องเตรียมก่อนเขียนข้อกำหนดผลิตภัณฑ์

ขั้นที่ 5 การพัฒนาผลิตภัณฑ์ (Product development)

(1) **การพัฒนาผลิตภัณฑ์ต้นแบบ (Prototype product development)** เมื่อได้แนวความคิดผลิตภัณฑ์และข้อกำหนดผลิตภัณฑ์แล้ว ขั้นตอนต่อไปนักพัฒนาผลิตภัณฑ์ต้องทำการพัฒนาผลิตภัณฑ์ต้นแบบโดยอาศัยการทำงานที่เป็นระบบและมีแบบแผน (Systematic experimentation) กำหนดปัจจัยที่ต้องการศึกษาและวางแผนการทดลอง การทำงานอย่างมีระบบขั้นตอนนี้จะทำให้ลดทั้งเวลาและค่าใช้จ่ายที่ไม่จำเป็นออกไป ขั้นตอนการพัฒนาผลิตภัณฑ์ต้นแบบเช่นเดียวกับขั้นตอนอื่นๆ นั่นคือ นักพัฒนาผลิตภัณฑ์ต้องระบุปัญหาของโครงการก่อน แล้วจึงกำหนดวัตถุประสงค์ที่แสดงถึงลำดับขั้นตอนย่อยในการทำงานและผลลัพธ์ที่จะได้เมื่อเสร็จสิ้นในแต่ละขั้นตอนย่อย เช่น ศึกษาวัตถุดิบและสภาวะในการแปรรูป ศึกษาชนิดภาชนะบรรจุที่เหมาะสม พิจารณาเลือกแผนการทดลองที่เหมาะสม และดำเนินการทดลอง วิเคราะห์และสรุปผล นำผลที่ได้มาประเมินแนวทางการพัฒนาผลิตภัณฑ์ต้นแบบให้ดียิ่งขึ้นเพื่อตอบสนองต่อความต้องการของผู้บริโภค

(2) **การพัฒนาผลิตภัณฑ์ที่เหมาะสม** ข้อมูลจากการพัฒนาผลิตภัณฑ์ต้นแบบจะให้นักพัฒนาผลิตภัณฑ์ทราบถึงทิศทางที่ต้องปรับปรุงผลิตภัณฑ์ ในขั้นตอนต่อไปคือการพัฒนาผลิตภัณฑ์ที่เหมาะสม (Product optimization) ผลิตภัณฑ์ที่เหมาะสมไม่ได้หมายถึงผลิตภัณฑ์ที่ดีที่สุด แต่เป็นผลิตภัณฑ์ที่สอดคล้องกับเป้าหมาย วัตถุประสงค์ ข้อจำกัดของโครงการ และกลยุทธ์ของบริษัท

(3) **การทดสอบคุณภาพผลิตภัณฑ์** ในระหว่างขั้นตอนการพัฒนาผลิตภัณฑ์ต้นแบบและผลิตภัณฑ์ที่เหมาะสมจำเป็นต้องวิเคราะห์ค่าคุณภาพของผลิตภัณฑ์เพื่อให้แน่ใจว่าผลิตภัณฑ์เป็นไปตามข้อกำหนดผลิตภัณฑ์ที่ได้ระบุไว้หรือไม่ สิ่งสำคัญที่ต้องคำนึงถึงคือ ค่าพารามิเตอร์ที่เลือกมาวิเคราะห์สอดคล้องกับคุณภาพหรือคุณลักษณะสำคัญของผลิตภัณฑ์ที่ผู้บริโภคยอมรับหรือไม่ การวิเคราะห์ค่าคุณภาพผลิตภัณฑ์แบ่งออกได้เป็น 2 วิธี คือ วิธีการวิเคราะห์ค่าทางตรง (Objective measurement) ได้แก่ การวิเคราะห์ค่าทางกายภาพ เคมี และจุลินทรีย์ และวิธีการวิเคราะห์ค่าทางอ้อม (Subjective measurement) ได้แก่ การประเมินคุณภาพทางประสาทสัมผัส

ในการประเมินคุณภาพทางประสาทสัมผัส จะใช้แบบสอบถามหรือใบทดสอบเป็นเครื่องมือในการรวบรวมข้อมูล ทั้งนี้วิธีการประเมินคุณภาพทางประสาทสัมผัสแบ่งออกได้เป็น 3 ประเภท คือ วิธีการประเมินความแตกต่างของผลิตภัณฑ์ (Difference test), วิธีการวิเคราะห์ลักษณะทางประสาทสัมผัสของผลิตภัณฑ์ (Descriptive test) และ วิธีการประเมินการยอมรับของผลิตภัณฑ์ (Acceptance test) ซึ่งแต่ละประเภทยังประกอบไปด้วยหลายวิธี แต่ละวิธีใช้ใบทดสอบที่แตกต่างกันตามวัตถุประสงค์ของการทดสอบและใช้สเกลในการวัดค่าที่แตกต่างกัน ดังนั้นผู้วิจัยจะต้องมีความรู้ความเข้าใจในหลักการของแต่ละวิธีเพื่อจะได้เลือกใช้วิธีทางสถิติได้อย่างถูกต้องเหมาะสม (รายละเอียดชนิดและประเภทสเกลของข้อมูลอ่านเพิ่มเติมในบทที่ 2)

การวางแผนการตลาดที่นิยมใช้ในงานวิจัยและพัฒนาผลิตภัณฑ์ เช่น การวางแผนการตลาดแบบ สุ่มสมบูรณ์ (รายละเอียดในบทที่ 8), การวางแผนการตลาดแบบสุ่มในบล็อกสมบูรณ์ (รายละเอียดในบทที่ 9) และการวางแผนการตลาดแบบแฟกทอเรียล (รายละเอียดในบทที่ 10)

การวิเคราะห์ทางสถิติที่นิยมใช้ในงานวิจัยและพัฒนาผลิตภัณฑ์ เช่น การทดสอบสมมติฐานงานวิจัย (รายละเอียดในบทที่ 5), การเปรียบเทียบความแตกต่างของค่าเฉลี่ยหนึ่งกลุ่มตัวอย่างและสองกลุ่มตัวอย่าง (รายละเอียดในบทที่ 6), การเปรียบเทียบความแตกต่างของค่าเฉลี่ยมากกว่าสองกลุ่มตัวอย่าง (รายละเอียดใน บทที่ 7), วิธีการแสดงผลตอบสนองแบบโครงร่างพื้นผิว (Response surface methodology, RSM) และการ วิเคราะห์การถดถอยเชิงพหุ (Regression analysis)

ขั้นที่ 6 การวางแผนทางการตลาดและการนำผลิตภัณฑ์ออกสู่ตลาด

(1) **การวางแผนการผลิตและการตลาด** การวางแผนการตลาดจะประกอบไปด้วยการกำหนด กลุ่มเป้าหมาย, แนวความคิดผลิตภัณฑ์, การคำนวณต้นทุนและกำหนดการขาย, การกำหนดส่วนผสมทางการ ตลาด, การวางแผนการผลิต และการกำหนดงบประมาณทางการตลาด

(2) **การทดสอบตลาด** ผลิตภัณฑ์สุดท้ายเมื่อได้นำไปทดสอบการยอมรับของผู้บริโภคแล้ว ขั้นตอน ต่อไปคือนำผลิตภัณฑ์ที่ได้มาทดสอบตลาด เพื่อเป็นการทำให้แน่ใจได้ว่าหากนำผลิตภัณฑ์ของบริษัทออกสู่ ตลาดสามารถแข่งขันกับผลิตภัณฑ์ชนิดเดียวกันของบริษัทอื่นได้ การทดสอบตลาดจะทำให้ได้ข้อมูลทาง การตลาดเพื่อนำมาใช้ในการกำหนดกลยุทธ์ทางการตลาด เช่น ราคา การส่งเสริมการขาย การทดสอบตลาด สามารถทำได้ 2 แบบคือ การทดสอบในสถานการณ์จริง (Real test marketing) และ การทดสอบใน สถานการณ์จำลอง (Simulated test marketing) ข้อมูลที่ได้จากการทดสอบตลาดจะนำไปใช้ในการพยากรณ์ ยอดขายหรือส่วนครองตลาดสำหรับตลาดอื่นๆ ได้ วิธีนี้เป็นที่นิยมใช้สำหรับทดสอบศักยภาพของผลิตภัณฑ์ใหม่

(3) **การนำผลิตภัณฑ์ออกสู่ตลาด** วิธีการนำผลิตภัณฑ์ออกสู่ตลาดแบ่งออกได้เป็น 3 ประเภท ได้แก่ 1) การนำผลิตภัณฑ์ออกวางตลาดทั่วประเทศพร้อมๆ กัน (National launch) มักใช้กับผลิตภัณฑ์ที่บริษัท มั่นใจมากกว่าจะประสบความสำเร็จโดยพิจารณาจากการทดสอบการยอมรับ, 2) การวางตลาดในบางพื้นที่ก่อน (Area launch) มักใช้กับผลิตภัณฑ์ที่ยังไม่แน่ใจว่าจะประสบความสำเร็จหรือไม่ หรือบริษัทขนาดเล็กที่ระบบ การกระจายสินค้ายังไม่ดีพอ และ 3) การขยายพื้นที่การวางจำหน่ายผลิตภัณฑ์ (Rolling launch) เมื่อได้ ทดลองนำผลิตภัณฑ์ออกสู่ตลาดแล้วพบว่าประสบความสำเร็จ จึงทำการขยายพื้นที่ในการวางจำหน่ายเพิ่มขึ้น

(4) **การประเมินผลการนำผลิตภัณฑ์ออกสู่ตลาด** โดยแบ่งออกเป็น 2 ขั้นตอน คือ 1) การ ประเมินผล กำหนดตัวชี้วัดและวิธีการประเมินให้ถูกต้องเพื่อประเมินความสำเร็จของผลิตภัณฑ์ โดยนำผลที่ได้ จากการประเมินได้จริงไปเปรียบเทียบกับจากการคาดคะเนหรือพยากรณ์ไว้ และ 2) การติดตามผล กำหนด

แผนการติดตามผลการปฏิบัติงาน เช่น การทำงานของฝ่ายขาย ปัญหาที่เกิดขึ้นคืออะไร และแนวทางแก้ไขที่นำไปใช้ก่อให้เกิดประสิทธิผลที่มากขึ้นหรือไม่

ขั้นที่ 7 การจัดการวงจรชีวิตผลิตภัณฑ์

วงจรชีวิตผลิตภัณฑ์แสดงถึงการเปลี่ยนแปลงยอดขายของผลิตภัณฑ์ตั้งแต่เริ่มนำผลิตภัณฑ์ออกสู่ตลาดจนกระทั่งผลิตภัณฑ์มียอดขายตกต่ำจนต้องนำออกจากตลาด ผลิตภัณฑ์แต่ละตัวมีวงจรชีวิตผลิตภัณฑ์ที่แตกต่างกัน ดังได้อธิบายในภาพที่ 1.2 ซึ่งแสดงกราฟวงจรชีวิตของผลิตภัณฑ์ จะเห็นว่ายอดขายและกำไรจะเกิดขึ้นจากความต้องการของตลาดที่มีต่อผลิตภัณฑ์ เมื่อผลิตภัณฑ์เติบโตเต็มที่และเริ่มตกต่ำ ยอดขายและกำไรจะลดลง ดังนั้นการนำเสนอผลิตภัณฑ์ใหม่ในเวลาที่เหมาะสมจึงเป็นเรื่องที่บริษัทจะต้องพิจารณาและคำนึงถึง

การจัดการวงจรชีวิตผลิตภัณฑ์ไม่ได้หมายถึงแค่การจัดการผลิตภัณฑ์ในช่วงถดถอยที่มียอดขายลดลงเท่านั้น แต่ยังหมายถึงการจัดการตลอดวงจรชีวิตผลิตภัณฑ์ การจัดการวงจรชีวิตผลิตภัณฑ์สามารถทำได้โดยอาศัยแนวทางของการจัดการโลจิสติกส์ (Logistics) เข้ามาช่วยในการจัดการประสานความร่วมมือระหว่างผู้จัดหา ผู้ผลิต ผู้จัดการคลังสินค้า ผู้กระจายสินค้า ผู้ค้าส่ง และร้านค้าปลีกต่างๆ ดังแสดงในตารางที่ 1.1

ตารางที่ 1.1 การจัดการวงจรชีวิตผลิตภัณฑ์ตามแนวทางโลจิสติกส์

วงจรชีวิตผลิตภัณฑ์	กิจกรรมที่เกิดขึ้นในแต่ละช่วง	แนวทางการจัดการโลจิสติกส์
ช่วงแนะนำผลิตภัณฑ์	มีการออกแบบผลิตภัณฑ์และนำสินค้าออกสู่ตลาด	การจัดการวัตถุดิบเป็นสิ่งสำคัญ ควรมีการประสานงานระหว่างผู้จัดหาวัตถุดิบและผู้ผลิต
ช่วงเจริญเติบโต	ผลิตภัณฑ์มียอดขายมากขึ้น เริ่มมีคู่แข่งเข้ามาแย่งส่วนแบ่งตลาด	การกระจายสินค้าเป็นสิ่งสำคัญ อาจมีการปรับปรุงสินค้าตามข้อมูลที่ได้จากลูกค้า ยังคงต้องมีการประสานงานระหว่างผู้จัดหาวัตถุดิบและผู้ผลิต
ช่วงอิ่มตัว	ยอดขายเติบโตเต็มที่ มีคู่แข่งเข้ามาแย่งส่วนแบ่งตลาดมากขึ้น	การแข่งขันเรื่องราคาเป็นสิ่งสำคัญ ต้นทุนต้องต่ำ เลือกผู้จัดหาวัตถุดิบที่มีราคาถูก รักษาการบริการ
ช่วงถดถอย	ปริมาณยอดขายลดลงจนต้องออกจากตลาดไป	การจัดการเก็บผลิตภัณฑ์จากลูกค้า กำจัดหรือนำไปปรับปรุงใหม่ ถือเป็นสิ่งสำคัญ

1.6 บทบาทและความสำคัญของสถิติในงานวิจัยและพัฒนา

สถิติ เป็นศาสตร์ที่ประกอบด้วย การเก็บรวบรวมข้อมูล การวิเคราะห์ข้อมูล การสรุปข้อมูล และการนำเสนอผลการวิเคราะห์เพื่อจะนำผลสรุปไปใช้ในการตัดสินใจด้านต่างๆ สถิติมีบทบาทในงานวิจัยเกือบทุกขั้นตอน โดยเริ่มตั้งแต่การวางแผนงานวิจัย การออกแบบงานวิจัย การกำหนดประชากรเป้าหมาย เทคนิคหรือแผนการเลือกตัวอย่าง การคำนวณขนาดตัวอย่าง การวิเคราะห์ผลที่ได้จากการรวบรวมข้อมูล การสรุปผล การหาสาเหตุหรือปัจจัย การเปรียบเทียบสิ่งต่างๆ การสรุปผล การนำเสนอผลงานวิจัย และการนำผลงานวิจัยไปประยุกต์ใช้

สถิติมีความสำคัญอย่างยิ่งต่อการวิจัย บางครั้งมีผู้วิจัยที่ใช้สถิติไม่ถูกต้องทำให้สรุปผลผิดพลาดซึ่งก่อให้เกิดความเสียหายต่องานที่เกี่ยวข้อง ปัจจุบันแม้เราจะเห็นว่าผู้รับจ้างวิเคราะห์ทางสถิติแต่จากที่ผู้เขียนได้กล่าวข้างต้นว่า “การทำวิจัยต้องมีความรู้ความเข้าใจในศาสตร์ของงานวิจัยนั้นควบคู่กับสถิติ” จึงมักพบในหลายครั้งว่าผู้รับจ้างวิเคราะห์เหล่านั้นใช้สถิติแบบไม่ถูกต้อง นอกจากนี้ในปัจจุบันการที่โปรแกรมทางสถิติได้ถูกพัฒนาไปอย่างรวดเร็วเพื่อให้สะดวกกับการใช้งานของผู้ใช้ แม้ว่าผู้ใช้งานจะป้อนข้อมูลและใช้คำสั่งวิเคราะห์ไม่ถูกต้องแต่โปรแกรมทางสถิติยังคงสามารถวิเคราะห์ผลทางสถิติให้โดยที่ผู้ใช้งานไม่ทราบว่าเป็นการวิเคราะห์ที่ไม่ถูกต้องจึงส่งผลให้สรุปผลงานวิจัยผิดพลาดได้

ตัวอย่างกรณีศึกษา ความเสียหายจากการใช้สถิติแบบไม่ถูกต้อง

โรงงานผลิตอาหารสัตว์แห่งหนึ่งต้องการซื้อเครื่องมือใหม่เพื่อนำมาใช้วิเคราะห์ไขมันในวัตถุดิบ เนื่องจากเครื่องมือเดิมใช้เวลาในการวิเคราะห์นาน ตัวอย่างละ 20 ชั่วโมง แต่เครื่องมือใหม่ใช้เวลาในการวิเคราะห์เพียงตัวอย่างละ 2 ชั่วโมง และเนื่องจากเครื่องมือใหม่มีราคา 4 ล้านบาท ดังนั้นก่อนโรงงานจะตัดสินใจซื้อได้ทำการสุ่มตัวอย่างวัตถุดิบมา 1 กระสอบเพื่อวิเคราะห์ไขมันด้วยเครื่องมือเดิมและเครื่องมือใหม่และเปรียบเทียบปริมาณไขมันที่วิเคราะห์ได้ ทั้งนี้โรงงานจะซื้อเครื่องมือใหม่หากพบว่าปริมาณไขมันที่วิเคราะห์ได้จากทั้งสองเครื่องมือแตกต่างกันอย่างมีนัยสำคัญทางสถิติ

เจ้าหน้าที่ห้องปฏิบัติการได้นำค่าปริมาณไขมันที่วิเคราะห์ได้จากทั้งสองเครื่องมือมาเปรียบเทียบความแตกต่างด้วยวิธี t-test ผลการวิเคราะห์พบว่าค่าเฉลี่ยปริมาณไขมันจากเครื่องมือใหม่แตกต่างกับเครื่องมือเดิมที่ระดับนัยสำคัญ 0.05 แต่เจ้าหน้าที่สรุปผลผิดว่าไม่แตกต่างกัน โรงงานจึงตัดสินใจซื้อเครื่องมือใหม่ เมื่อโรงงานสั่งซื้อเครื่องมือใหม่มาใช้แล้วมาพบภายหลังว่าได้ค่าไขมันที่ไม่ถูกต้อง สร้างความเสียหายให้แก่โรงงานเป็นอย่างมาก เนื่องจากฝ่ายผลิตต้องนำค่าปริมาณไขมันไปใช้ในการคำนวณสูตรอาหารสัตว์ ทำให้อาหารสัตว์ที่ผลิตออกมาไม่ได้คุณภาพตามที่กำหนดไว้ ลูกค้าส่งคืนและฟ้องร้องเรียกค่าเสียหาย จากตัวอย่างนี้แสดงให้เห็นว่า การใช้สถิติที่ไม่ถูกต้องทำให้การสรุปผลและตัดสินใจผิดพลาดอันส่งผลให้เกิดความเสียหายต่อองค์กรได้

ประโยชน์ของสถิติในงานวิจัยและพัฒนา มีดังนี้

(1) ช่วยวางแผนงานวิจัย ตั้งแต่การเก็บรวบรวมข้อมูล การกำหนดขนาดตัวอย่าง การออกแบบงานวิจัยในขั้นตอนต่างๆ

สถิติเข้ามาเกี่ยวข้องกับงานวิจัยเกือบทุกขั้นตอน โดยเริ่มตั้งแต่การวางแผนงานวิจัย การออกแบบงานวิจัย การกำหนดประชากรเป้าหมาย (Target population) การเลือกตัวอย่าง การคำนวณขนาดตัวอย่าง ผู้วิจัยต้องกำหนดว่าประชากรเป้าหมายเป็นใครที่จะตอบวัตถุประสงค์ของงานวิจัยได้ เช่น ถ้าต้องการพัฒนาผลิตภัณฑ์สำหรับผู้สูงอายุ ในระหว่างขั้นตอนการพัฒนาสูตรควรใช้ผู้ทดสอบที่เป็นผู้สูงอายุในการทดสอบ ประเมินความชอบผลิตภัณฑ์ หากใช้กลุ่มวัยรุ่นซึ่งมีลักษณะการรับรู้ลักษณะทางประสาทสัมผัสที่แตกต่างกับผู้สูงอายุจะทำให้ได้ลักษณะของผลิตภัณฑ์ที่ไม่สามารถตอบสนองความต้องการของผู้สูงอายุได้ ทั้งนี้ผู้สูงอายุมีปัจจัยที่ต้องคำนึงถึงหลายปัจจัย เช่น ความสามารถในการเคี้ยวและกลืน การรับรู้รส ดังนั้นนักวิจัยต้องกำหนดเทคนิคการเลือกตัวอย่างซึ่งเป็นตัวแทนของประชากรซึ่งมีหลายวิธี เช่น การเลือกตัวอย่างที่ใช้ความน่าจะเป็น (Probability sampling) และการเลือกตัวอย่างที่ไม่ใช้ความน่าจะเป็น (Non-probability sampling) ดังแสดงในตารางที่ 1.2

ตารางที่ 1.2 วิธีการเลือกตัวอย่างที่ใช้และไม่ใช้ความน่าจะเป็น

การเลือกตัวอย่างที่ใช้ความน่าจะเป็น (Probability sampling)	การเลือกตัวอย่างที่ไม่ใช้ความน่าจะเป็น (Non-probability sampling)
การสุ่มตัวอย่างแบบง่าย (Simple random sampling)	การเลือกตัวอย่างแบบสะดวก (Convenience sampling)
การสุ่มตัวอย่างแบบเป็นระบบ (Systematic sampling)	การเลือกตัวอย่างแบบเจาะจง (Purposive sampling)
การเลือกตัวอย่างแบบแบ่งชั้นภูมิ (Stratified sampling)	การเลือกตัวอย่างแบบบังเอิญ (Accidental sampling)
การเลือกตัวอย่างแบบกลุ่ม (Cluster sampling)	การเลือกตัวอย่างแบบโควตา (Quota sampling)
การเลือกตัวอย่างแบบหลายขั้นตอน (Multi-stage sampling)	การเลือกตัวอย่างแบบบอกต่อ (Snowball sampling)

การกำหนดขนาดตัวอย่าง (Sample size determination) ผู้วิจัยมักจะมีคำถามว่าจะใช้ขนาดตัวอย่างเท่าใดจึงจะเหมาะสม ขนาดตัวอย่างขึ้นกับ เทคนิคการเลือกตัวอย่าง งบประมาณ กำลังคน เวลานอกจากนี้ยังขึ้นอยู่กับ ขนาดประชากร ความคลาดเคลื่อนจากการใช้ตัวอย่างแทนประชากร หรือความผิดพลาดสูงสุดที่ยอมให้เกิดในการประมาณค่าพารามิเตอร์ และความผันแปรของข้อมูล

(2) เป็นเครื่องมือที่ช่วยให้เห็นเหตุการณ์หรือความเป็นจริงที่เกิดขึ้น เช่น การลวกเปลือกแผ่นก่อนทอดช่วยให้ลดการอมน้ำมันลงได้, ผู้ทดสอบให้คะแนนความชอบเค้กสูตร A มากกว่าสูตร B และ ผู้บริโภคเพศหญิงให้คะแนนความชอบผลิตภัณฑ์มากกว่าเพศชาย กรณีที่ผู้วิจัยใช้ตัวอย่างเป็นตัวแทนประชากร เมื่อต้องการสรุปลักษณะที่สำคัญของประชากรจะใช้วิธีการทดสอบสมมติฐาน

เครื่องมือในการเก็บรวบรวมข้อมูลอาจจะเป็นเครื่องชั่งน้ำหนัก เครื่องวิเคราะห์คุณภาพด้านเนื้อสัมผัส เครื่องวิเคราะห์คุณภาพด้านสี แบบสอบถาม ในงานวิจัยและพัฒนาผลิตภัณฑ์นั้นนอกจากการตรวจสอบคุณภาพวัตถุดิบและผลิตภัณฑ์โดยใช้การวิเคราะห์ค่าด้วยเครื่องมือ เช่น การวิเคราะห์ค่าคุณภาพทางกายภาพทางเคมี และทางจุลินทรีย์แล้ว ยังเกี่ยวข้องกับการใช้คนหรือผู้ทดสอบในการประเมินคุณภาพทางประสาทสัมผัสซึ่งใช้แบบสอบถามหรือใบทดสอบเป็นเครื่องมือในการรวบรวมข้อมูล

(3) หาสาเหตุที่ทำให้เกิดเหตุการณ์ต่างๆ เมื่อนักวิจัยสามารถหาสาเหตุของปัญหาได้แล้วจะนำไปสู่การวางแผนหรือหาทางแก้ไข เช่น นักวิจัยพบว่าสาเหตุที่ทำให้มะละกอแฉิมอบแห้งมีสีคล้ำเนื่องจากเอนไซม์ Polyphenol oxidase (PPO) ที่มีอยู่ในมะละกอ ดังนั้นก่อนการแฉิมจึงแก้ไขโดยนำชิ้นมะละกอไปลวกในน้ำร้อนเพื่อยับยั้งการทำงานของเอนไซม์ที่ก่อให้เกิดปฏิกิริยาสีน้ำตาล

กรณีที่วัตถุประสงค์ของงานวิจัย คือ การหาสาเหตุหรือศึกษาหาปัจจัยที่ส่งผลต่อเรื่องใดเรื่องหนึ่ง เช่น การศึกษาพฤติกรรมและความต้องการของผู้บริโภค การศึกษาปัจจัยที่มีผลต่อการตัดสินใจซื้อสินค้าของผู้บริโภค จะใช้เทคนิคการวิเคราะห์ความสัมพันธ์ โดยทดสอบสมมติฐานเพื่อหาความสัมพันธ์

(4) ช่วยในการประเมินผลงาน เช่น ประเมินผลความพึงพอใจของลูกค้า และ ประเมินผลการยอมรับของผู้บริโภคที่มีต่อผลิตภัณฑ์ที่พัฒนา และช่วยในการตัดสินใจ เช่น การคัดเลือกวัตถุดิบที่เหมาะสม การตัดสินใจนำผลิตภัณฑ์ออกสู่ตลาดเมื่อพิจารณาข้อมูลจากการทดสอบการยอมรับของผู้บริโภค

(5) ใช้ในการตรวจสอบความเชื่อถือได้ของข้อมูลที่เก็บรวบรวม การตรวจสอบความเชื่อถือของเครื่องมือที่ใช้ในการเก็บรวบรวมข้อมูลทำได้ 3 ประเภท คือ

- ความเชื่อถือได้ที่วัดจากความเสถียรภาพ ได้แก่ การทดสอบความสอดคล้องของคนกลุ่มเดียวกันแต่เวลาต่างกัน เช่น ใช้แบบสอบถามคำถามเดิมกับคนกลุ่มเดิมแต่ต่างเวลากัน ถ้าคำตอบ 2 ครั้งเหมือนกันหรือสอดคล้องกันจะเรียกว่ามีความน่าเชื่อถือ

- ความเชื่อถือได้ที่วัดจากเครื่องมือที่สามารถทดแทนกันได้ เป็นการวัดความเชื่อถือได้ของเครื่องมือโดยการสอบถามหรือวัดจากคนๆ เดียวกันหรือสิ่งของเดียวกัน โดยใช้เครื่องมือต่างกัน เช่น ใช้แบบสอบถาม 2 ชุด (มีลักษณะของคำถามเหมือนกัน) สอบถามคนๆ เดียวกัน ถ้าข้อมูลจากแบบสอบถาม 2 ชุดสอดคล้องกันแสดงว่ามีความเชื่อถือได้

- ความเชื่อถือได้ที่วัดความสอดคล้องภายในชุดเดียวกัน ในงานวิจัยบางครั้งเป็นการยากที่จะทดสอบความเชื่อถือได้จากคนๆ เดิมในเวลาที่แตกต่างกัน หรือให้คนๆ เดียวทำแบบสอบถาม 2 ชุด ดังนั้นวิธีการนี้

จะทดสอบความเชื่อถือกับคนอื่นๆ ในครั้งเดียวเลย เช่น ให้ผู้ทดสอบตอบแบบสอบถามโดยแบบสอบถามแบ่งออกเป็น 2 ส่วนเพื่อทดสอบความเชื่อถือได้ในครั้งเดียว

ในงานวิจัยและพัฒนาผลิตภัณฑ์ทางด้านอุตสาหกรรมอาหาร โดยส่วนใหญ่การนำสถิติไปประยุกต์ใช้จะเกี่ยวข้องกับขั้นตอนการพัฒนาผลิตภัณฑ์ เช่น การหาสูตรและการหาสภาวะที่เหมาะสมในกระบวนการผลิต (Optimization) การพัฒนาสูตรใหม่ๆ การจัดการและควบคุมคุณภาพของผลิตภัณฑ์ การพัฒนาผลิตภัณฑ์ใหม่ หรือการปรับปรุงคุณภาพผลิตภัณฑ์ให้เป็นที่ยอมรับของผู้บริโภค งานวิจัยและพัฒนาเหล่านี้เป็นสิ่งที่ท้าทายสำหรับนักวิจัยและพัฒนาผลิตภัณฑ์ ตัวอย่างคำถามที่มักเกิดขึ้นกับนักพัฒนาผลิตภัณฑ์ เช่น



- ถ้าเปลี่ยนแปลงสภาวะในการผลิตจะส่งผลกระทบต่อคุณภาพของผลิตภัณฑ์อย่างไร ?
- สูตรหรือวัตถุดิบตัวไหนที่ควรเลือกใช้เพื่อให้ได้ผลิตภัณฑ์ที่ผู้บริโภคยอมรับ ?
- ต้องใช้สภาวะในการผลิตหรือวัตถุดิบตัวไหนจะทำให้ลดต้นทุนกระบวนการผลิต ?
- ถ้าเปลี่ยนวัตถุดิบใหม่จะมีผลต่อการยอมรับของผู้บริโภคหรือไม่ ?
- ทำอย่างไรจะทำให้ผลิตภัณฑ์ยังคงอยู่ในตลาดได้ ต้องปรับปรุงอย่างไร ?
- ปัจจัยอะไรที่ส่งผลกระทบต่ออายุการเก็บของผลิตภัณฑ์ ?
- ทำอย่างไรจะทำให้กระบวนการผลิตได้ผลผลิตที่สูงขึ้น ?

การตอบคำถามดังกล่าวต้องอาศัยกระบวนการทำวิจัยที่เป็นขั้นตอนและมีการวางแผนงานอย่างเป็นระบบ โดยใช้สถิติเป็นเครื่องมือในการตัดสินใจ จากที่กล่าวไว้ในตอนต้นของหนังสือบทนี้การที่นักวิจัยจะทำวิจัยให้ประสบความสำเร็จได้นั้นต้องมีความรู้ความเข้าใจในศาสตร์ของงานวิจัยนั้นควบคู่ไปกับความรู้ทางสถิติ เช่น การวิจัยหาสภาวะที่เหมาะสมในกระบวนการผลิตนอกจากนักวิจัยจะต้องมีความรู้ด้านวัตถุดิบ กระบวนการผลิต และการวิเคราะห์ค่าคุณภาพยังต้องมีความรู้ทางสถิติด้วย

ในหนังสือเล่มนี้ได้อธิบายการวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรมสำเร็จรูปทางสถิติ SPSS โดยมุ่งเน้นการใช้สถิติในขั้นตอนการพัฒนาผลิตภัณฑ์เป็นหลัก อาทิเช่น การคัดเลือกสูตรพื้นฐาน การศึกษาผลของวัตถุดิบต่อคุณภาพผลิตภัณฑ์ การศึกษาผลของสภาวะการผลิตต่อคุณภาพผลิตภัณฑ์ นอกจากนี้ยังได้อธิบายถึงสถิติเชิงพรรณนาที่นิยมใช้ในขั้นตอนการสร้างแนวความคิดผลิตภัณฑ์และการทดสอบตลาด



บทที่ 2

สถิติพื้นฐานและการทบทวน

สถิติเป็นเครื่องมือที่นักวิจัยใช้ในการแสวงหาคำตอบและหาข้อสรุปในสิ่งที่ต้องการศึกษาด้วยวิธีการที่เป็นที่ยอมรับและใช้กันอย่างกว้างขวางในเชิงวิชาการไม่ว่าจะเป็นการใช้สถิติเพื่อพรรณนาคุณสมบัติของสิ่งที่สนใจ หรือการใช้สถิติเพื่อศึกษาหาความแตกต่างระหว่างกลุ่ม หรือการใช้สถิติเพื่อศึกษาหาความสัมพันธ์ระหว่างตัวแปร นอกจากนี้ผู้วิจัยจะต้องเลือกใช้วิธีการทางสถิติในการวางแผนการทดลองและวิเคราะห์ข้อมูลให้ถูกต้องแล้ว ผู้วิจัยยังต้องรู้หลักการเบื้องต้นของสถิติต่างๆ เช่น ประเภทของสถิติ ชนิดและประเภทของข้อมูล คำศัพท์ต่างๆ ที่ใช้ในงานวิจัย การสุ่มตัวอย่างและการทำซ้ำการทดลอง ซึ่งนักวิจัยควรทำความเข้าใจคำศัพท์ที่เกี่ยวข้องเหล่านี้เพื่อให้สามารถประยุกต์ใช้ในการวางแผนการทดลองสำหรับงานวิจัยด้านพัฒนาผลิตภัณฑ์ได้อย่างถูกต้อง นักวิจัยควรคำนึงอยู่เสมอว่า ชนิดของข้อมูลที่แตกต่างกันย่อมใช้วิธีการทางสถิติในการวิเคราะห์ที่แตกต่างกัน หากไม่เข้าใจในชนิดของข้อมูลอาจส่งผลให้เลือกวิธีการวิเคราะห์ทางสถิติที่ผิดและท้ายที่สุดคือการสรุปผลที่ผิดพลาดและคลาดเคลื่อนไปได้

2.1 ประเภทของสถิติสำหรับงานวิจัย

สถิติ หมายถึง การวางแผนสำหรับการเก็บรวบรวมข้อมูล การเก็บรวบรวมข้อมูล การนำเสนอข้อมูล เบื้องต้น การวิเคราะห์ข้อมูล และการสรุปผล ซึ่งสามารถนำผลสรุปนั้นมาช่วยในการตัดสินใจ สถิติสามารถแบ่งได้หลายประเภทตามเกณฑ์จำแนกต่างๆ เช่น แบ่งตามเกณฑ์จำนวนตัวแปรตาม (สถิติขั้นพื้นฐาน และ สถิติขั้นสูง), แบ่งตามเกณฑ์ลักษณะวิชา (สถิติเชิงทฤษฎี และ สถิติประยุกต์), แบ่งตามข้อกำหนดเกี่ยวกับกลุ่มประชากร (สถิติพารามетริก และ สถิตินอนพารามетริก) และแบ่งตามเกณฑ์วิธีวิเคราะห์ข้อมูล (สถิติเชิงพรรณนา และ สถิติอนุมาน) ซึ่งในหนังสือเล่มนี้จะอธิบายสถิติที่แบ่งตามเกณฑ์วิธีวิเคราะห์ข้อมูล ดังนี้

2.1.1 สถิติเชิงพรรณนา (Descriptive statistics)

สถิติเชิงพรรณนาเป็นสถิติเบื้องต้นที่สุดที่ใช้ในการสรุปลักษณะสำคัญของกลุ่มประชากรหรือกลุ่มตัวอย่าง ซึ่งไม่ได้สรุปอ้างอิงไปยังกลุ่มอื่น เช่น อาจารย์ผู้สอนวิชาฟิสิกส์พบว่าคะแนนสอบของนิสิตหมู่ 1 มีค่าเฉลี่ยเท่ากับ 35.65 อาจารย์จะสรุปอ้างอิงไปยังนิสิตหมู่ 2 ว่ามีคะแนนสอบเฉลี่ยเท่ากับ 35.65 ไม่ได้ วิธีวิเคราะห์ข้อมูลสถิติเชิงพรรณนาประกอบด้วย ขนาดตัวอย่าง, ความถี่, ร้อยละ, ค่ากลาง (เช่น ค่าเฉลี่ย

ค่ามัธยฐาน และ ค่าฐานนิยม), ค่าการกระจาย (เช่น ค่าแปรปรวน ค่าส่วนเบี่ยงเบนมาตรฐาน ค่าเปอร์เซ็นต์ไทล์ และ ค่าควอไทล์) และ กราฟต่างๆ (เช่น กราฟแท่ง, กราฟ Histogram, กราฟ Boxplot และ กราฟวงกลม)

2.1.2 สถิติอนุมาน (Inferential statistics)

สถิติอนุมานเป็นสถิติที่มุ่งศึกษากับกลุ่มตัวอย่างแล้วสรุปผลที่ได้ศึกษาไปยังกลุ่มประชากร โดยอาศัยทฤษฎีความน่าจะเป็น เช่น โรงงานน้ำตาลต้องการวัดค่าความหวานของอ้อยที่เกษตรกรนำมาขายเพื่อใช้ในการประเมินค่าตอบแทน ถ้ารถแต่ละคันมีอ้อยประมาณ 1 ตัน ดังนั้นถ้าใช้สถิติอนุมานโรงงานจะทำการสุ่มตัวอย่างอ้อยจากรถแต่ละคันในจำนวนที่เหมาะสมเพื่อนำมาวัดค่าความหวาน ไม่จำเป็นต้องวัดอ้อยทุกลำในรถคันนั้น

สถิติอนุมานสามารถแบ่งออกได้เป็น 2 ประเภท คือ

(1) **สถิติประมาณ (Estimation statistics)** เป็นสถิติที่ใช้ประมาณค่าพารามิเตอร์ (Parameter) ด้วยค่าทางสถิติ (Statistics) เช่น ประมาณค่าเฉลี่ย μ ของประชากร ด้วยค่าสถิติ $\bar{X}$ ของกลุ่มตัวอย่าง ความหมายของค่าพารามิเตอร์และค่าสถิติ มีดังนี้

- ค่าพารามิเตอร์ คือ ค่าที่คำนวณได้จากข้อมูลประชากรจึงเป็นค่าที่บรรยายหรือแสดงลักษณะประชากร เช่น ค่าเฉลี่ย (μ) น้ำหนักสุกรทุกตัวในฟาร์ม A

- ค่าสถิติ คือ ค่าที่คำนวณได้จากกลุ่มตัวอย่าง จึงเป็นค่าที่บรรยายหรือแสดงลักษณะของกลุ่มตัวอย่าง เช่น ค่าเฉลี่ย ($\bar{X}$) น้ำหนักสุกร 50 ตัวที่สุ่มตรวจจากสุกรทั้งหมดในฟาร์ม A

(2) **สถิติทดสอบ (Test statistics)** เป็นสถิติเกี่ยวกับการทดสอบสมมติฐาน ซึ่งแบ่งพิจารณาออกได้เป็น 2 กลุ่ม คือ การสรุปลักษณะประชากร และ การหาความสัมพันธ์

- การสรุปลักษณะประชากร กรณีที่งานวิจัยใช้ข้อมูลกลุ่มตัวอย่างจะต้องใช้การทดสอบสมมติฐานทางสถิติเพื่ออ้างอิงลักษณะประชากร เช่น การทดสอบเกี่ยวกับค่าสัดส่วนประชากร การทดสอบเกี่ยวกับค่าเฉลี่ยประชากร

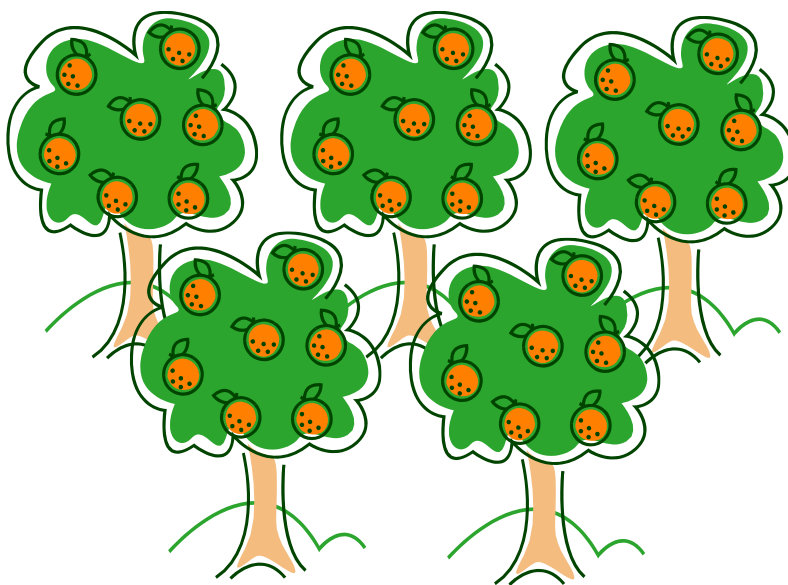
- การหาความสัมพันธ์ เป็นการหาสาเหตุหรือในงานวิจัยมักใช้คำว่า “การศึกษาปัจจัยที่ส่งผลต่อ....” ทั้งนี้สถิติที่ใช้ในการหาความสัมพันธ์มีหลายวิธี ซึ่งหลักการเลือกใช้ขึ้นอยู่กับวัตถุประสงค์ของการทำวิจัย และ ชนิดของสเกลข้อมูล (ซึ่งจะกล่าวถึงในข้อ 2.4.3) เทคนิคที่ใช้หาความสัมพันธ์ เช่น การวิเคราะห์ z-test และ t-test, การวิเคราะห์ความแปรปรวน (Analysis of Variance: ANOVA), การวิเคราะห์ Chi-square test, การวิเคราะห์ความถดถอย (Regression analysis), การวิเคราะห์สหสัมพันธ์ (Correlation analysis), การวิเคราะห์จำแนกกลุ่ม (Discriminant analysis), การวิเคราะห์ปัจจัย (Factor analysis), การแบ่งกลุ่ม (Cluster analysis) และ การวิเคราะห์สถิติไม่พารามิเตอร์ (Nonparametric test)

2.2 ประชากร (Population) และตัวอย่าง (Sample)

2.2.1 ประชากร (Population) หมายถึง ทุกหน่วยในเรื่องที่สนใจศึกษา คำว่า หน่วย อาจหมายถึง คน สัตว์ สิ่งของ พืช และ องค์กร เช่น นักวิจัยสนใจทดสอบความหวานของส้มในไร่ A ซึ่งมีต้นส้มอยู่ 100 ต้น แบ่งออกเป็น 5 แปลง แปลงละ 20 ต้น ดังนั้น ส้มทุกลูกในไร่ส้ม คือ ประชากร และ ส้ม 1 ลูก คือ 1 หน่วย

2.2.2 ตัวอย่าง (Sample) หมายถึง ส่วนย่อยหรือบางส่วนของประชากร เช่น นักวิจัยสุ่มเก็บส้ม 5 แปลง แปลงละ 5 ต้น ต้นละ 5 ลูก กรณีนี้จะได้ส้มจำนวน 125 ลูก ซึ่งใช้เป็นตัวอย่างหรือตัวแทนของประชากร ส้มทั้งไร่ การกำหนดขนาดตัวอย่างนั้นขึ้นอยู่กับหลายปัจจัย เช่น เทคนิคการเลือกตัวอย่าง งบประมาณ กำลังคน เวลา ขนาดประชากร ความแปรผันของข้อมูล

กรณีที่ประชากรมีขนาดใหญ่ ผู้วิจัยอาจไม่สามารถเก็บข้อมูลจากทุกหน่วยของประชากรได้ จะต้องเก็บข้อมูลจากตัวอย่างซึ่งถือว่าตัวอย่างเป็นตัวแทนของประชากร ดังนั้น ตัวอย่างจะต้องมีลักษณะเหมือนประชากร ถ้ากลุ่มตัวอย่างมีลักษณะแตกต่างจากประชากรแล้ว ข้อมูลตัวอย่างที่ได้ย่อมไม่น่าเชื่อถือและไม่ควรนำมาใช้ในงานวิจัย ภาพที่ 2.1 แสดงการเปรียบเทียบประชากรและตัวอย่างส้ม หากเจ้าของไร่ส้มทำการวัดความหวานส้มทุกลูกเพื่อจะมั่นใจในคุณภาพส้มก็จะมีส้มพอให้จำหน่ายได้ ดังนั้นเจ้าของไร่ส้มต้องวางแผนการสุ่มเก็บตัวอย่างส้มเพื่อให้แน่ใจว่าจะเป็นตัวแทนคุณภาพของส้มทั้งไร่



ภาพที่ 2.1 ประชากร คือ ส้มทุกลูก และ ตัวอย่าง คือ ส้มบางส่วนที่ถูกสุ่มเก็บมาเป็นตัวแทนประชากร

2.3 ตัวแปร

ตัวแปร (Variable) หมายถึง คุณลักษณะต่างๆ ของบุคคล สิ่งของ หรือ สภาพแวดล้อม ทั้งในเชิงปริมาณ และเชิงคุณภาพ ซึ่งมีค่าแตกต่างกันไป เช่น เพศ อายุ ส่วนสูง ทั้งนี้สามารถจำแนกตัวแปรได้หลายเกณฑ์ ดังนี้

2.3.1 จำแนกตามความเป็นเหตุและผล

(1) **ตัวแปรอิสระ หรือ ตัวแปรต้น (Independent variable)** คือ ตัวแปรต้นเหตุที่ทำให้เกิดผล ต่อ ตัวแปรที่สนใจทำการการศึกษา (ตัวแปรตาม)

(2) **ตัวแปรตาม (Dependent variable)** คือ ตัวแปรที่สนใจทำการการศึกษา

(3) **ตัวแปรแทรกซ้อน (Extraneous variable)** คือ ตัวแปรที่มีผลต่อตัวแปรตาม แต่ไม่ต้องการ ศึกษา

(4) **ตัวแปรปรับ หรือ ตัวแปรกำกับ (Moderator variable)** คือ ตัวแปรที่มีปฏิสัมพันธ์กับตัวแปรอิสระที่จะมีผลต่อตัวแปรตาม

(5) **ตัวแปรส่งผ่าน หรือ ตัวแปรคั่นกลาง (Mediator variable)** คือ ตัวแปรที่เป็นผลของตัวแปรอิสระและเป็นเหตุของตัวแปรตาม

2.3.2 จำแนกตามลักษณะตัวแปร

(1) **ตัวแปรเชิงคุณภาพ (Qualitative variable)** คือ ตัวแปรที่มีค่าต่างๆ อยู่ในรูปคุณภาพหรือคุณลักษณะ (Attribute) ไม่สามารถบอกปริมาณได้ เช่น เพศ ศาสนา วิธีการทอด ชนิดแป้ง

(2) **ตัวแปรเชิงปริมาณ (Quantitative variable)** คือ ตัวแปรที่มีค่าต่างๆ อยู่ในรูปจำนวนหรือขนาด สามารถบอกปริมาณความมากน้อยได้ เช่น น้ำหนัก อายุ ระดับความหวาน ปริมาณสารทดแทนไขมัน ปริมาณโปรตีน คะแนนความชอบ

2.3.3 จำแนกตามความต่อเนื่อง

(1) **ตัวแปรต่อเนื่อง (Continuous variable)** คือ ตัวแปรที่มีค่าต่างๆ วัดได้ละเอียดเป็นไปได้ทุกค่า เช่น น้ำหนัก อายุ ปริมาณไขมัน ค่าความแข็ง

(2) **ตัวแปรไม่ต่อเนื่อง (Discrete variable)** คือ ตัวแปรที่มีค่าต่างๆ แยกจากกันโดยเด็ดขาด ไม่สามารถวัดได้อย่างละเอียดทุกค่า เช่น จำนวนนักเรียน สามารถวัดได้ 1, 2, 3,... 15 คน แต่บอกไม่ได้ว่ามีจำนวนนักเรียน 1.1, 1.2, 1.3 คน

2.4 ประเภทของข้อมูล

ข้อมูล หมายถึง ข้อเท็จจริงหรือความจริงที่เกิดขึ้นซึ่งอาจจะเป็นตัวเลขหรือข้อความ หรือประกอบด้วย ข้อมูลตัวเลขและข้อความ เช่น

“ภาควิชาพัฒนาผลิตภัณฑ์อยู่ในคณะอุตสาหกรรมเกษตร” จัดเป็นข้อมูลที่อยู่ในรูปแบบข้อความ

“สำนักงานเศรษฐกิจการเกษตรรายงานผลผลิตข้าวนาปี ในปีเพาะปลูก 2559/60 รวมทั้งประเทศ มีเนื้อที่เพาะปลูก 58,645,474 ไร่” จำนวนเนื้อที่เพาะปลูก จัดเป็นข้อมูลที่อยู่ในรูปแบบตัวเลข

“ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเกษตรศาสตร์ เป็นที่แรกในประเทศไทยที่เปิดการเรียนการสอนในหลักสูตรพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร ปัจจุบันมีคณาจารย์ จำนวน 17 คน เป็นผู้ชาย 2 คน และผู้หญิง 15 คน” จัดเป็นข้อมูลที่อยู่ในรูปแบบตัวเลขและข้อความ

การแปลความหมายและการประมวลผลข้อมูลอาจได้มาจากการสังเกต การรวบรวม การวัด การทดลอง ในการทำวิจัยสิ่งสำคัญที่ต้องคำนึงถึง คือ ประเภทของข้อมูล เนื่องจากใช้เป็นตัวกำหนดวิธีการวิเคราะห์ทางสถิติ การจำแนกประเภทของข้อมูลสามารถจำแนกได้หลายลักษณะ เช่น จำแนกตามแหล่งที่มาของข้อมูล จำแนกตามลักษณะของข้อมูล และ จำแนกตามสเกลของข้อมูล

2.4.1 จำแนกตามแหล่งที่มาของข้อมูล ออกได้เป็น 2 ชนิด ดังนี้

(1) **ข้อมูลปฐมภูมิ (Primary data)** เป็นข้อมูลที่ผู้วิจัยรวบรวมจากแหล่งข้อมูลโดยตรง อาจเป็นการรวบรวมโดยใช้วิธี สัมภาษณ์ ทำการทดลอง และ สังเกตการณ์ หรืออาจเป็นข้อมูลที่มีผู้เก็บรวบรวมไว้แต่ยังไม่มี การจัดการอย่างเป็นระบบ หมวดยุ่ การที่ผู้วิจัยทำการเก็บข้อมูลปฐมภูมิเองมีข้อดีคือ ผู้วิจัยได้ข้อมูลที่ครบตามที่ตนเองต้องการ เป็นข้อมูลที่ไม่ลำสมัย และสามารถมั่นใจในขั้นตอนการเก็บข้อมูล ส่วนข้อเสีย คือ มีค่าใช้จ่ายสูงและสิ้นเปลืองเวลาในการเก็บข้อมูล ในงานพัฒนาผลิตภัณฑ์นั้นการสำรวจพฤติกรรมและความต้องการของผู้บริโภคที่มีต่อผลิตภัณฑ์เป็นสิ่งที่นักวิจัยควรทำการสำรวจด้วยตัวเองเนื่องจากในปัจจุบันพฤติกรรมและความต้องการของผู้บริโภคเปลี่ยนแปลงอย่างรวดเร็วซึ่งเป็นสาเหตุมาจากสื่อออนไลน์และสภาพสังคม การใช้ข้อมูลเดิมที่มีผู้ทำวิจัยไว้แล้วอาจไม่สามารถใช้อ้างอิงได้ทั้งหมด

(2) **ข้อมูลทุติยภูมิ (Secondary data)** เป็นข้อมูลที่ผู้อื่นหรือหน่วยงานอื่นได้ทำการเก็บรวบรวมไว้แล้วและมักเป็นข้อมูลที่ได้อัวิเคราะห์ในขั้นต้นแล้ว อาจแสดงในรูปแบบยอดรวม จำนวน ค่าเฉลี่ย และร้อยละ เช่น เว็บไซต์ของสำนักงานเศรษฐกิจการเกษตรรายงานผลผลิตข้าวนาปี ในปีเพาะปลูก 2562/63 รวมทั้งประเทศ มีเนื้อที่เพาะปลูก 61,197,134 ไร่, เนื้อที่เก็บเกี่ยว 54,108,276 ไร่ และผลผลิต 24,064,170 ตัน ตามลำดับ นักวิจัยสามารถนำข้อมูลทุติยภูมิมาใช้ประกอบการทำวิจัยซึ่งมีข้อดีคือ สะดวก ประหยัดทั้งเวลาและค่าใช้จ่าย แต่ข้อเสียคือผู้วิจัยอาจได้ข้อมูลไม่ครบตามที่ตัวเองต้องการ ข้อมูลอาจล้าสมัย และไม่สามารถทราบวิธีการเก็บข้อมูลได้อย่างละเอียดว่าน่าเชื่อถือเพียงใด ผู้วิจัยจึงต้องพิจารณาอย่างระมัดระวังในการเลือกใช้

ข้อมูลทุติยภูมิ ทั้งนี้ในงานพัฒนาผลิตภัณฑ์ผู้วิจัยควรทำการสืบค้นข้อมูลทุติยภูมิก่อนเริ่มทำการวิจัยเพื่อเป็นแนวทางในการสร้างแนวคิดผลิตภัณฑ์และการตัดสินใจริเริ่มโครงการวิจัย แหล่งของข้อมูลทุติยภูมิ เช่น ฐานข้อมูลสิทธิบัตรและอนุสิทธิบัตร รายงานการวิจัย วารสารงานวิจัย และ ฐานข้อมูลออนไลน์

2.4.2 จำแนกตามลักษณะข้อมูล ออกได้เป็น 2 ชนิด ดังนี้

(1) **ข้อมูลเชิงคุณภาพ (Qualitative data)** บางครั้งเรียกว่า ข้อมูลเชิงกลุ่ม หรือข้อมูลจำแนกกลุ่ม เป็นข้อมูลที่อยู่ในรูปข้อความและแสดงลักษณะที่แตกต่างกัน เช่น

“เพศ” หมายถึง แบ่งคนเป็น 2 กลุ่มตามเพศสภาพที่ต่างกัน คือ ชาย และ หญิง

“ประเภทนิสิต” เช่น แบ่งนิสิตเป็น 3 กลุ่มตามระดับการศึกษา ได้แก่ นิสิตปริญญาตรี นิสิตปริญญาโท และ นิสิตปริญญาเอก

“เกรด” เช่น แบ่งเกรดออกเป็น 5 กลุ่ม ได้แก่ A, B, C, D และ F

“สถานภาพนิสิต” เช่น แบ่งนิสิตออกเป็น 2 กลุ่มตามเกรดคะแนนเฉลี่ยสะสม ได้แก่ นิสิตรอปินิจ (มีคะแนนเฉลี่ยสะสม น้อยกว่า 2.00) และนิสิตปกติ (มีคะแนนเฉลี่ยสะสม มากกว่าหรือเท่ากับ 2.00)

“เกรดคุณภาพน้ำปลา” เช่น แบ่งน้ำปลาออกเป็น 2 กลุ่มตามคุณภาพ ได้แก่ น้ำปลาแท้ และ น้ำปลาผสม

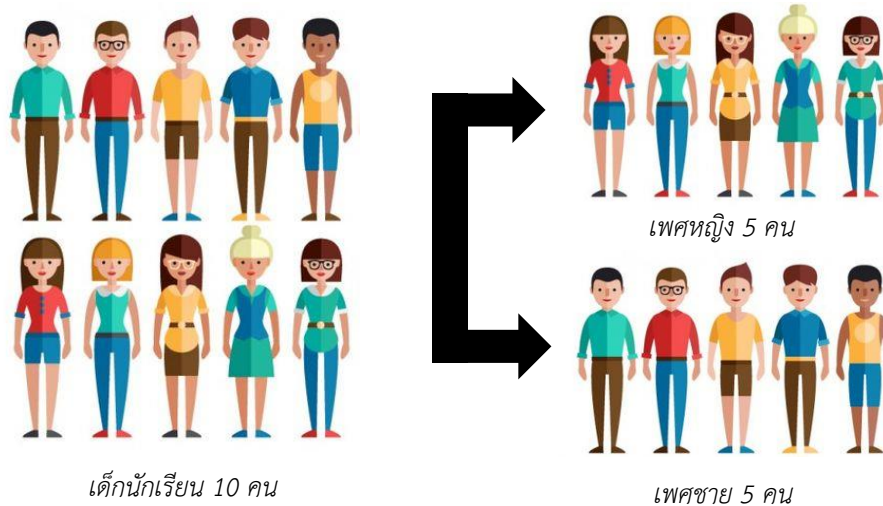
“การยอมรับผลิตภัณฑ์” เช่น แบ่งการยอมรับผลิตภัณฑ์เป็น 2 กลุ่ม ได้แก่ ยอมรับ และ ไม่ยอมรับ

“กลุ่มอายุ” เช่น แบ่งกลุ่มอายุผู้บริโภคมอบออกเป็น 2 กลุ่ม ได้แก่ กลุ่มวัยเรียน และ กลุ่มวัยทำงาน กรณีนี้เราสามารถวัดความแตกต่างได้ระดับหนึ่ง นั่นคือ กลุ่มวัยทำงานมีอายุมากกว่ากลุ่มวัยเรียน แม้จะไม่สามารถทราบปริมาณความแตกต่างที่แน่นอน

(2) **ข้อมูลเชิงปริมาณ (Quantitative data)** เป็นข้อมูลตัวเลขที่สามารถวัดค่าได้ว่ามีค่ามากหรือน้อย เช่น รายได้ อายุ เกรดเฉลี่ยสะสม ยอดขาย ปริมาณไขมัน ปริมาณโปรตีน ราคาขาย ระดับความพึงพอใจ และ คะแนนความชอบ

2.4.3 จำแนกตามสเกลของข้อมูล ออกได้เป็น 4 ชนิด ดังนี้

(1) **สเกลนามกำหนด หรือ สเกลนามบัญญัติ (Nominal scale)** เป็นสเกลที่แสดงถึงการแบ่งกลุ่มหรือจำแนกกลุ่ม โดยไม่สามารถบอกได้ว่ากลุ่มไหนมากกว่าหรือน้อยกว่า จึงถือว่าเป็นสเกลที่หยาบที่สุด ดังแสดงในภาพที่ 2.2 เป็นการแบ่งกลุ่มนักเรียนจำนวน 10 คน ออกเป็น 2 กลุ่มตามเพศสภาพที่แตกต่าง คือ เพศชาย และ เพศหญิง



ภาพที่ 2.2 การแบ่งกลุ่มโดยใช้สเกลนามกำหนดแบ่งนักเรียนออกเป็น 2 กลุ่ม คือ เพศชาย และ หญิง

ในการวิเคราะห์ค่าทางสถิติจะต้องกำหนดรหัสที่เป็นตัวเลขให้แต่ละกลุ่มเพื่อให้สะดวกในการสร้างแฟ้มข้อมูล เช่น ข้อมูลด้านเพศ กำหนดให้รหัส 1 คือ เพศชาย และรหัส 2 คือ เพศหญิง, ข้อมูลประเภทนิสิต กำหนดให้รหัส 1 คือ นิสิตระดับปริญญาตรี รหัส 2 คือ นิสิตระดับปริญญาโท และรหัส 3 คือ นิสิตระดับปริญญาเอก ผู้วิจัยต้องจำไว้เสมอว่ารหัสตัวเลขที่กำหนดนี้เป็นเพียงการแบ่งหรือจำแนกกลุ่ม ไม่สามารถนำมาใช้บอกความมากน้อยได้ ในการวิเคราะห์ทางสถิติสามารถนำข้อมูลที่ได้จากสเกลนามกำหนด มาวิเคราะห์หา ความถี่ ร้อยละ และ รฐานนิยม ได้

(2) **สเกลอันดับ (Ordinal scale)** เป็นสเกลที่สามารถใช้ในการแบ่งกลุ่มเช่นเดียวกับสเกลนามกำหนด แต่จะให้รายละเอียดมากกว่าโดยแสดงถึงความแตกต่างระหว่างกลุ่มได้ สามารถจัดอันดับได้ว่ากลุ่มไหนมากกว่าหรือน้อยกว่า ดีกว่าหรือแย่กว่า จากภาพที่ 2.3 แสดงการเปรียบเทียบระดับความสูงของเด็กนักเรียน 3 คน ซึ่งเราสามารถเรียงลำดับความสูงจากมากไปน้อยได้ บอกความแตกต่างได้แต่ไม่สามารถบอกขนาดความแตกต่างได้ กรณีความสูงของเด็กนักเรียนทั้งสามคน ความแตกต่างของลำดับที่ 1 กับ 2 ไม่จำเป็นต้องเท่ากับความแตกต่างของลำดับที่ 2 และ 3



ภาพที่ 2.3 การเปรียบเทียบระดับความสูงของเด็กนักเรียน 3 คน

ตัวอย่างสเกลอันดับที่มักใช้ในการสร้างแบบสอบถาม เช่น อายุที่แบ่งเป็นช่วง (เช่น 15-20 ปี, 21-25 ปี และ 26-30 ปี) รายได้ต่อเดือน (เช่น น้อยกว่า 5,000 บาท, 5,000-10,000 บาท และ มากกว่า 10,000 บาท ขึ้นไป) และ ระดับการศึกษา (เช่น มัธยมศึกษา, ปริญญาตรี, ปริญญาโท และ ปริญญาเอก)

ในกรณีที่ผู้ตอบแบบสอบถามคนที่ 1 มีอายุในช่วง 15-20 ปี และมีรายได้ น้อยกว่า 5,000 บาท และ ผู้ตอบแบบสอบถามคนที่ 2 มีอายุในช่วง 21-25 ปี และมีรายได้ 5,000-10,000 บาท

จากข้อมูลดังกล่าว นักวิจัยสามารถทราบได้ว่า ผู้ตอบแบบสอบถามคนที่ 1 มีอายุและรายได้น้อยกว่า ผู้ตอบแบบสอบถามคนที่ 2 แต่ไม่ทราบอายุและรายได้ที่แท้จริงของทั้งสองคน

ในการจัดอันดับหรือเรียงลำดับสามารถเรียงจากมากไปน้อย หรือ น้อยไปมากก็ได้ ในการวิเคราะห์ทางสถิติสามารถนำข้อมูลที่ได้จากสเกลอันดับมาวิเคราะห์หา ความถี่ ร้อยละ ฐานนิยม และ ค่ามัธยฐาน ได้

(3) สเกลอันดับหรือ สเกลแบบช่วง (Interval scale) เป็นสเกลที่ให้รายละเอียดได้มากกว่า สเกลนามกำหนดและสเกลอันดับ สามารถวัดขนาดความแตกต่างได้ บอกได้ว่ามากกว่าหรือน้อยกว่า ดีกว่าหรือ แย่กว่า พอใจมากกว่าหรือน้อยกว่า นิยมใช้สเกลนี้ในงานวิจัยที่เกี่ยวกับการศึกษาพฤติกรรมและความต้องการของผู้บริโภค สำนวจความคิดเห็น ประเมินระดับความพึงพอใจ สเกลอันดับนี้สามารถระบุระยะห่างของแต่ละสเกลด้วยช่วงที่เท่าๆ กัน แต่เป็นสเกลที่ไม่มีค่าศูนย์ที่แท้จริง (เป็นค่าศูนย์สมมติเท่านั้น) นั่นคือ จุดเริ่มต้นหรือค่าเริ่มต้นไม่มีความหมาย และจุดเริ่มต้นอาจจะเท่ากันหรือไม่เท่ากันก็ได้ ตัวอย่างสเกลอันดับ เช่น ระดับความพึงพอใจของนิสิตต่อการเรียน แบ่งเป็น 5 ระดับ ได้แก่ 1 = น้อยที่สุด, 2 = น้อย, 3 = ปานกลาง, 4 = มาก และ 5 = มากที่สุด ในการวิเคราะห์ทางสถิติสามารถนำข้อมูลที่ได้จากสเกลอันดับมาวิเคราะห์ค่าสถิติได้หลายแบบ เช่น ค่าเฉลี่ย ค่ามัธยฐาน ค่าส่วนเบี่ยงเบนมาตรฐาน และ เปอร์เซ็นไทล์

(4) สเกลอัตราส่วน (Ratio scale) เป็นสเกลที่ให้ข้อมูลได้มากที่สุด กล่าวคือ เป็นข้อมูลที่สามารถระบุขนาดได้ จึงทำให้สามารถเปรียบเทียบความแตกต่างได้ สเกลนี้มีค่าศูนย์ที่แท้จริง หรือ จุดเริ่มต้นเป็นค่าที่มีความหมาย เช่น ยอดขาย น้ำหนัก ส่วนสูง อายุ ปริมาณสินค้า ปริมาณน้ำตาล และ ปริมาณโปรตีน ในการวิเคราะห์ทางสถิติสามารถนำข้อมูลที่ได้จากสเกลอัตราส่วนมาวิเคราะห์ค่าสถิติได้หลายแบบ เช่น ค่าเฉลี่ย ค่ามัธยฐาน ค่าส่วนเบี่ยงเบนมาตรฐาน และ เปอร์เซ็นไทล์

ในการทำวิจัยนั้นข้อมูลที่น่ามาใช้อาจเป็นข้อมูลปฐมภูมิและข้อมูลทุติยภูมิ และอาจประกอบด้วย ข้อมูลเชิงคุณภาพและเชิงปริมาณ หรืออาจจะประกอบด้วยสเกลข้อมูลทั้ง 4 สเกล ผู้วิจัยควรพิจารณาข้อมูลที่ต้องการวิเคราะห์ทางสถิติว่าเป็นข้อมูลชนิดใดเพื่อจะได้เลือกใช้วิธีวิเคราะห์ทางสถิติให้เหมาะสม เช่น กรณีถ้าผู้วิจัยต้องการสอบถามอายุของผู้ตอบแบบสอบถาม ถ้าเป็นคำถามปลายเปิดให้ผู้ตอบแบบสอบถามระบุอายุที่แท้จริง เช่น อายุ.....ปี ข้อมูลเช่นนี้ถือเป็นข้อมูลเชิงปริมาณที่ใช้สเกลอัตราส่วน แต่ถ้านักวิจัยให้ผู้ตอบแบบสอบถามเลือกคำตอบที่มีลักษณะเป็นช่วงอายุ เช่น 20-25 ปี, 26-30 ปี, 31-35 ปี และ 36-40 ปี ข้อมูล

เช่นนี้ถือเป็นข้อมูลเชิงคุณภาพที่ใช้สเกลอันดับ จากกรณีดังกล่าวจะเห็นได้ว่าแม้ผู้วิจัยสนใจสอบถามเรื่องอายุเหมือนกันแต่ใช้คำถามที่แตกต่างกัน ชนิดของข้อมูลที่ได้ก็จะแตกต่างกันด้วย ดังนั้นวิธีการวิเคราะห์ทางสถิติก็จะแตกต่างกัน ภาพที่ 2.4 แสดงการเปรียบเทียบความแตกต่างของสเกลทั้ง 4 ชนิด และตารางที่ 2.1 แสดงการสรุปวิธีวิเคราะห์ข้อมูลที่ใช้กับข้อมูลโดยจำแนกตามชนิดสเกลที่ใช้

สเกลนามกำหนด (Nominal scale)	<table><tr><td>A</td><td>B</td><td>C</td><td>D</td><td>E</td></tr><tr><td>1</td><td>2</td><td>3</td><td>4</td><td>5</td></tr><tr><td>●</td><td>■</td><td>▲</td><td>◆</td><td>♥</td></tr></table>	A	B	C	D	E	1	2	3	4	5	●	■	▲	◆	♥
A	B	C	D	E												
1	2	3	4	5												
●	■	▲	◆	♥												
สเกลอันดับ (Ordinal scale)																
สเกลอันตรภาค (Interval scale)																
สเกลอัตราส่วน (Ratio scale)																

ภาพที่ 2.4 เปรียบเทียบความแตกต่างของสเกลทั้ง 4 ชนิด

ตารางที่ 2.1 สรุปวิธีวิเคราะห์ข้อมูลที่ใช้กับข้อมูลที่มีสเกลต่างกัน

ชนิดข้อมูล	ลักษณะข้อมูล	ตัวอย่างวิธีวิเคราะห์ทางสถิติ
สเกลนามกำหนด หรือ สเกลนามบัญญัติ (Nominal scale)	เป็นสเกลที่ใช้ในการแบ่งหรือจำแนก กลุ่มได้อย่างเดียว เช่น <i>เพศ (ชาย, หญิง)</i> <i>ความชอบ (ชอบ, ไม่ชอบ)</i> <i>ระดับการศึกษา (ป.ตรี, โท, เอก)</i> <i>อาชีพ (ราชการ, รัฐวิสาหกิจ, นิสิต)</i>	ความถี่, ร้อยละ, ฐานนิยม, Chi-square test, Binomial test
สเกลอันดับ (Ordinal scale)	เป็นสเกลที่ใช้ในการจัดเรียงลำดับ ข้อมูลสามารถบอกความมากและน้อย แต่ไม่สามารถบอกขนาดความ แตกต่างได้ เช่น <i>เมื่อเรียงลำดับความสูงจากมากไป น้อย พบว่านาย A, B, C มีความสูง อันดับที่ 1, 2, 3 กรณีนี้ไม่สามารถ บอกได้ว่าแต่ละคนมีความสูงแตกต่าง กันเท่าไร</i>	ค่ามัธยฐาน, ฐานนิยม, เปอร์เซ็นไทล์, สหสัมพันธ์ลำดับที่ (Rank - order correlation), Binomial test, Nonmetric multidimensional scaling test
สเกลอันตรภาค หรือ สเกลแบบช่วง (Interval scale)	เป็นสเกลที่ใช้บอกขนาดความ แตกต่างได้ แต่สเกลนี้ไม่มีค่าศูนย์ สัมบูรณ์ (Absolute zero) หรือ ไม่มี ค่าศูนย์ที่แท้จริง เช่น <i>ระดับความพึงพอใจ 1 = น้อย, 2 = ปานกลาง, 3 = มาก</i>	ค่าเฉลี่ย, ค่าส่วนเบี่ยงเบนมาตรฐาน, การทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ย, Correlation analysis, Discriminant analysis, Regression analysis, Analysis of variance, Metric multidimensional scaling test
สเกลอัตราส่วน (Ratio scale)	เป็นสเกลที่ใช้บอกขนาดความ แตกต่างได้ สเกลนี้มีค่าศูนย์สัมบูรณ์ และจุดเริ่มต้นมีความหมาย เช่น <i>ความยาว, ความสูง, ปริมาณโปรตีน</i>	ค่าเฉลี่ย, ค่าส่วนเบี่ยงเบนมาตรฐาน, การทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ย, Correlation analysis, Discriminant analysis, Regression analysis, Analysis of variance, Factor analysis, Cluster analysis, Metric multidimensional scaling test ** สามารถใช้สถิติได้เกือบทุกวิธีในการ วิเคราะห์ข้อมูล

2.5 การทดลอง (Experiment)

การทดลอง คือ วิธีการตรวจสอบเพื่อให้ได้มาซึ่งข้อเท็จจริงใหม่ หรือ เพื่อยืนยันหรือปฏิเสธผลการทดลองที่ผ่านมา โดยการทดลองแบ่งออกได้เป็น 3 ประเภท ได้แก่

(1) **การทดลองเบื้องต้น (Preliminary experiment)** เป็นการทดลองเบื้องต้นเพื่อนำผลที่ได้ไปศึกษาในการทดลองครั้งต่อไป

(2) **การทดลองขั้นวิกฤติ (Critical experiment)** เป็นการทดลองที่เกี่ยวข้องกับการศึกษาเปรียบเทียบความแตกต่างระหว่างอิทธิพลของสิ่งทดลองที่มีต่อค่าตอบสนอง เป็นการทดลองที่ต้องการหาคำตอบอย่างชัดเจนเพื่อนำผลการศึกษาไปใช้งานต่อไป

(3) **การทดลองขั้นสาธิต (Demonstration experiment)** เป็นการทดลองที่ต้องการแสดงการเปรียบเทียบให้เห็นถึงความแตกต่างระหว่างอิทธิพลของสิ่งทดลองใหม่กับสิ่งทดลองเดิมที่ใช้อยู่

2.6 ขั้นตอนในการวางแผนการทดลอง

(1) **กำหนดปัญหาที่ต้องการศึกษา** ผู้วิจัยต้องตั้งปัญหาที่สนใจศึกษาและกำหนดขอบเขตของการวิจัยให้ชัดเจน รวมทั้งลักษณะข้อมูลที่ต้องการเก็บรวบรวม ตลอดจนขั้นตอนการเก็บรวบรวมข้อมูล

(2) **กำหนดวัตถุประสงค์ของการทดลอง** กำหนดสมมติฐานที่ต้องการทดสอบ การกำหนดวัตถุประสงค์ของงานวิจัยที่ชัดเจน จะทำให้ผู้วิจัยทราบว่าสิ่งที่ต้องการทราบจริงๆ คืออะไร วัตถุประสงค์ต้องมีความชัดเจนได้ใจความ อาจแยกย่อยเป็นข้อๆ ได้โดยเรียงลำดับความสำคัญหรือลำดับการศึกษาวิจัยในแต่ละขั้น ในงานวิจัยด้านพัฒนาผลิตภัณฑ์มักเริ่มต้นด้วยการสำรวจพฤติกรรมและความต้องการของผู้บริโภคก่อนการพัฒนาสูตรและกรรมวิธีการผลิต ยกตัวอย่างการกำหนดวัตถุประสงค์ของโครงการวิจัย เช่น 1) เพื่อสำรวจพฤติกรรมและความต้องการของผู้บริโภคที่มีต่อการบริโภคน้ำผลไม้ 2) เพื่อพัฒนาสูตรและกรรมวิธีการผลิตน้ำผลไม้เสริมเส้นใย และ 3) เพื่อทดสอบการยอมรับของผู้บริโภคต่อน้ำผลไม้เสริมเส้นใย

(3) **การตัดสินใจเลือกใช้แผนการทดลองที่เหมาะสม** การวางแผนการทดลองถือเป็นขั้นตอนที่สำคัญมาก เนื่องจากการวางแผนการทดลองมีมากมายหลายชนิด แต่ละชนิดมีความเหมาะสมต่อปัญหาที่จะศึกษาแตกต่างกัน การเลือกใช้แผนการทดลองชนิดใดนั้นนอกจากจะขึ้นอยู่กับวัตถุประสงค์ของการศึกษาและจำนวนสาเหตุของความคลาดเคลื่อนในการทดลองแล้ว ยังขึ้นอยู่กับทรัพยากรในการทำวิจัย เช่น แรงงาน ทุน และเวลาอีกด้วย ในงานพัฒนาผลิตภัณฑ์ใช้การวางแผนการทดลองด้วยหลายวัตถุประสงค์ เช่น เพื่อพัฒนาสูตรเพื่อคัดเลือกสูตรพื้นฐาน เพื่อเปรียบเทียบกับตัวอย่างควบคุมหรือตัวอย่างมาตรฐาน เพื่อศึกษาพฤติกรรมและความต้องการของผู้บริโภค เพื่อศึกษาอายุการเก็บ และเพื่อศึกษาสภาวะที่เหมาะสมในการผลิต

2.7 สิ่งทดลอง (Treatment)

สิ่งทดลอง หรือ ทรีทเมนต์ หมายถึง ลักษณะประจำตัวของหน่วยทดลอง หรือวิธีการที่ใช้ปฏิบัติต่อหน่วยทดลอง ถ้านักวิจัยต้องการเปรียบเทียบชนิดของไขมันที่มีผลต่อความนุ่มของเค้ก ชนิดของไขมัน คือ ลักษณะประจำตัวของไขมัน เช่น เนย มาการีน และ น้ำมันพืช ในกรณีนี้การทดลองจะมี 3 สิ่งทดลอง ได้แก่ เค้กที่ใช้เนย เค้กที่ใช้มาการีน และเค้กที่ใช้ไขมันพืช หากนักวิจัยต้องการศึกษาผลของปริมาณสารทดแทนไขมันที่มีผลต่อความนุ่มของเค้ก กรณีนี้ปริมาณสารทดแทนไขมันที่ใช้ในการผลิตเค้กที่ระดับต่างๆ คือ วิธีปฏิบัติต่อหน่วยทดลอง สิ่งทดลองอาจเกิดจากปัจจัยเดียว เช่น ต้องการศึกษาค่าผลของอุณหภูมิในการทอดต่อการอมน้ำมันของข้าวเกรียบทอด หรือสิ่งทดลองอาจเกิดจากหลายปัจจัย เช่น ต้องการศึกษาค่าผลของอุณหภูมิและเวลาในการทอดต่อการอมน้ำมันของข้าวเกรียบทอด

ในงานวิจัยด้านพัฒนาผลิตภัณฑ์นั้นสิ่งทดลองอาจเป็นได้ทั้งลักษณะประจำตัวของหน่วยทดลอง และวิธีการที่ใช้ปฏิบัติต่อหน่วยทดลอง ดังนี้

(1) **ลักษณะประจำตัวของหน่วยทดลอง** เช่น การเปรียบเทียบความถี่ในการไปดูหนังของกลุ่มผู้ทดสอบเพศชายและเพศหญิง (เพศ คือ ลักษณะประจำตัวของหน่วยทดลอง), พฤติกรรมการซื้อขนมขบเคี้ยวของผู้บริโภคกลุ่มวัยเรียนและกลุ่มวัยทำงาน (อาชีพของกลุ่มผู้บริโภค คือ ลักษณะประจำตัวของหน่วยทดลอง), การเปรียบเทียบค่าความแข็งของขนมขบเคี้ยวที่ใช้แป้งต่างชนิดกัน (ชนิดแป้ง คือ ลักษณะประจำตัวของหน่วยทดลอง)

(2) **วิธีการที่ใช้ปฏิบัติต่อหน่วยทดลอง** เช่น การศึกษาผลของอุณหภูมิในการอบต่อความชื้นของมะเขือเทศแช่อบแห้ง (อุณหภูมิในการอบที่ระดับต่างๆ คือ วิธีการที่ใช้ปฏิบัติต่อมะเขือเทศแช่อบแห้ง), การศึกษาผลของเวลาในการทอดต่อการอมน้ำมันของมันฝรั่งทอด (เวลาในการทอดที่ระดับต่างๆ คือ วิธีการที่ใช้ปฏิบัติต่อหน่วยมันฝรั่ง), การศึกษาผลของการแทนที่แป้งสาลีด้วยแป้งกล้วยดิบต่อปริมาณเส้นใยในขนมปัง (ปริมาณการแทนที่แป้งสาลีด้วยแป้งกล้วยที่ระดับต่างๆ คือ วิธีการที่ใช้ปฏิบัติต่อขนมปัง)

ในบางครั้งสิ่งทดลอง หรือ ทรีทเมนต์ อาจถูกเรียกตามวัตถุประสงค์ของการทดลอง เช่น **สิ่งทดลองควบคุม หรือ ทรีทเมนต์ควบคุม (Control treatment)** ยกตัวอย่างเช่น กรณีนักวิจัยต้องการศึกษาผลของปริมาณสารทดแทนไขมันต่อความนุ่มของเค้ก นักวิจัยจึงทำการเปรียบเทียบความนุ่มของเค้กที่ใช้สารทดแทนไขมันแทนที่เนย (ที่ระดับ 25 และ 50%) กับ เค้กที่ใช้เนย (ใช้สารทดแทนไขมัน 0%) เรียกเค้กที่ใช้เนยนี้ว่า สิ่งทดลองควบคุม ซึ่งในการทดลองนี้จะประกอบด้วย 3 สิ่งทดลองตามปริมาณสารทดแทนไขมันที่ใช้แทนที่เนย ได้แก่ เค้กที่ใช้สารทดแทนไขมัน 0%, เค้กที่ใช้สารทดแทนไขมัน 25% และ เค้กที่ใช้สารทดแทนไขมัน 50%

ในการวัดอิทธิพลหรือผลของสิ่งทดลอง นักวิจัยต้องทำการวัด **ผลตอบสนอง หรือ คำสังเกต (Response)** ของหน่วยตัวอย่างสุ่มซึ่งอาจจะเป็นส่วนหนึ่งของการทดลองหรืออาจจะเป็นหน่วยทดลองทุก

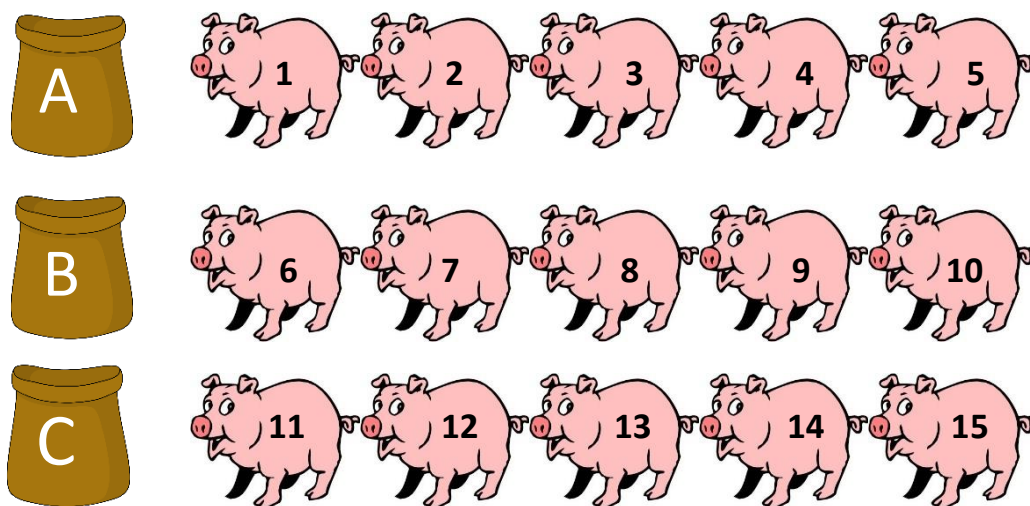
หน่วยก็ได้ เช่น ในการทำเค้ก 1 สูตร จะได้เค้ก 15 ชิ้น ผู้วิจัยอาจสุ่มมาวัดค่าความนุ่ม 5 ชิ้น หรือ วัดความนุ่ม ทั้ง 15 ชิ้นก็ได้ ค่าความนุ่มที่ได้เรียกว่า ผลตอบสนอง หรือ ค่าสังเกต

2.8 หน่วยทดลอง (Experimental Unit)

หน่วยทดลอง หมายถึง หน่วยหรือกลุ่มของวัตถุทดลองที่ได้รับสิ่งทดลองหรือทรีทเมนต์ชนิดเดียวกัน หน่วยทดลองอาจจะเป็นหน่วยทดลองเดี่ยว หรือ หน่วยทดลองกลุ่มก็ได้

(1) ตัวอย่างกรณีหน่วยทดลองเดี่ยว โรงงานแห่งหนึ่งพัฒนาอาหารเลี้ยงสุกร 3 สูตร (สูตร A, สูตร B และ สูตร C) และต้องการศึกษาว่าอาหารสูตรใดที่จะเพิ่มน้ำหนักสุกรได้ดีที่สุด จึงทำการสุ่มสุกรสายพันธุ์เดียวกันที่มีอายุและน้ำหนักเท่ากันจากคอกเดียวกัน จำนวน 15 ตัวมาใช้ในการทดลอง แบ่งสุกรออกเป็น 3 กลุ่ม กลุ่มละ 5 ตัว นำมาเลี้ยงด้วยอาหารสูตร A (สุกรตัวที่ 1-5) , สูตร B (สุกรตัวที่ 6-10) และสูตร C (สุกรตัวที่ 11-15) เมื่อครบ 1 เดือนจะทำการชั่งน้ำหนักสุกร

ในกรณีนี้จะมีจำนวนหน่วยทดลองเดี่ยวทั้งหมด เท่ากับ 15 หน่วยทดลอง (สุกร 15 ตัว) ดังภาพที่ 2.5 และสามารถเขียนตารางบันทึกผลน้ำหนักได้ดังตารางที่ 2.2



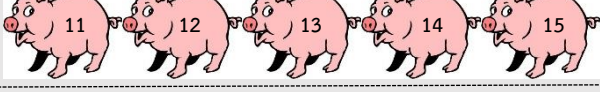
ภาพที่ 2.5 ตัวอย่างหน่วยทดลองเดี่ยว (ตัวเลขแสดงจำนวนหน่วยทดลอง)

ตารางที่ 2.2 น้ำหนักสุกรที่ได้รับอาหารสูตร A, B และ C เมื่อครบ 1 เดือน

สูตรอาหาร	น้ำหนักสุกร (กิโลกรัม)				
สูตร A	ตัวที่ 1	ตัวที่ 2	ตัวที่ 3	ตัวที่ 4	ตัวที่ 5
สูตร B	ตัวที่ 6	ตัวที่ 7	ตัวที่ 8	ตัวที่ 9	ตัวที่ 10
สูตร C	ตัวที่ 11	ตัวที่ 12	ตัวที่ 13	ตัวที่ 14	ตัวที่ 15

(2) ตัวอย่างกรณีหน่วยทดลองกลุ่ม โรงงานแห่งหนึ่งพัฒนาอาหารเลี้ยงสุกร 3 สูตร (สูตร A, สูตร B และสูตร C) และต้องการศึกษาว่าอาหารสูตรใดที่จะเพิ่มน้ำหนักสุกรได้ดีที่สุด จึงทำการสุ่มสุกรสายพันธุ์เดียวกันที่มีอายุและน้ำหนักเท่ากันจากคอกที่ 1 และคอกที่ 2 จำนวนคอกละ 15 ตัวมาใช้ในการทดลอง ในแต่ละคอกจะแบ่งสุกรออกเป็น 3 กลุ่ม กลุ่มละ 5 ตัว เพื่อนำมาเลี้ยงด้วยอาหารสูตร A, สูตร B และ สูตร C และเมื่อครบ 1 เดือนจะทำการชั่งน้ำหนักสุกร

ในกรณีนี้อาหารแต่ละสูตรจะประกอบด้วยหน่วยทดลองกลุ่ม 2 กลุ่ม (จากคอกที่ 1 และคอกที่ 2) โดยแต่ละกลุ่มมีหน่วยทดลองเดี่ยว 5 หน่วย (สุกร 5 ตัว) ดังนั้นอาหารแต่ละสูตรจะมีจำนวนหน่วยทดลองเดี่ยว 10 หน่วย ซึ่งสามารถสรุปได้ว่าการทดลองอาหารสุกรทั้ง 3 สูตรมีหน่วยทดลองกลุ่มทั้งหมด 6 กลุ่ม และหน่วยทดลองเดี่ยวทั้งหมด 30 หน่วยดังภาพที่ 2.6 และสามารถเขียนตารางบันทึกผลน้ำหนักได้ดังตารางที่ 2.3

		คอกที่ 1 หน่วยทดลองกลุ่มที่ 1
		คอกที่ 1 หน่วยทดลองกลุ่มที่ 2
		คอกที่ 1 หน่วยทดลองกลุ่มที่ 3
		คอกที่ 2 หน่วยทดลองกลุ่มที่ 4
		คอกที่ 2 หน่วยทดลองกลุ่มที่ 5
		คอกที่ 2 หน่วยทดลองกลุ่มที่ 6

ภาพที่ 2.6 ตัวอย่างหน่วยทดลองกลุ่ม (ตัวเลขแสดงจำนวนหน่วยทดลองเดี่ยว)

ตารางที่ 2.3 น้ำหนักสุกรที่ได้รับอาหารสูตร A, B และ C จากสุกร 2 คอก

สูตรอาหาร	สุกรคอกที่	หน่วยทดลองกลุ่มที่	น้ำหนักสุกร (กิโลกรัม)				
			หน่วยทดลองเดี่ยวที่ (สุกรตัวที่)				
สูตร A	1	1	1	2	3	4	5
สูตร B	1	2	6	7	8	9	10
สูตร C	1	3	11	12	13	14	15
สูตร A	2	4	16	17	18	19	20
สูตร B	2	5	21	22	23	24	25
สูตร C	2	6	26	27	28	29	30

2.9 ปัจจัยและระดับของปัจจัย (Factor and level of factor)

ปัจจัย (Factor) หมายถึง สิ่งที่น่าสนใจศึกษาว่ามีผลต่อหน่วยทดลองหรือผลตอบสนองหรือไม่ ระดับของปัจจัย (Level of factor) หมายถึง ความแตกต่างของปัจจัย อาจสามารถเรียงลำดับตามความมากน้อยได้ ปัจจัยที่มีระดับแตกต่างกันสามารถใช้เป็นสิ่งที่ทดลองหรือทรีทเมนต์ได้ ทั้งนี้ความแตกต่างระหว่าง “ปัจจัย” และ “สิ่งทดลอง” คือ สิ่งทดลองหมายถึงลักษณะประจำตัวของหน่วยทดลอง หรือวิธีที่ปฏิบัติต่อหน่วยทดลอง ซึ่งอาจประกอบด้วยปัจจัยเดียว หรือหลายปัจจัยรวมกันก็ได้ ถ้าทำการทดลองโดยใช้หลายปัจจัยร่วมกัน และแต่ละปัจจัยมีหลายระดับ จำนวนของสิ่งทดลองจะเกิดจากระดับของปัจจัยร่วมกัน เรียกว่า **การจัดกลุ่มสิ่งทดลอง (Treatment combination)** ซึ่งมีจำนวนเท่ากับ ผลคูณของระดับแต่ละปัจจัย

กรณีที่ 1 สิ่งทดลองหรือทรีทเมนต์ประกอบด้วยปัจจัยเดียว

นักวิจัยต้องการศึกษาว่าอุณหภูมิในการทอดมีผลต่อการอมน้ำมันของนั้กเกิดจากแป้งถั่วเหลืองหรือไม่ จึงทำการทอดนั้กเกิดจากแป้งถั่วเหลืองที่อุณหภูมิ 140, 150 และ 160 องศาเซลเซียส แล้วนำนั้กเกิดที่ได้ไปวิเคราะห์ปริมาณไขมันโดยวิธีทางเคมีในห้องปฏิบัติการ กรณีนี้เมื่อพิจารณาจะได้รายละเอียดของข้อมูล ดังนี้

- ปัจจัยที่ศึกษามี 1 ปัจจัย คือ อุณหภูมิในการทอด
- ระดับของปัจจัยที่ศึกษา มี 3 ระดับ คือ 140, 150 และ 160 องศาเซลเซียส
- จำนวนสิ่งทดลองมี 3 สิ่งทดลอง (เท่ากับจำนวนระดับของปัจจัย) ได้แก่
 - สิ่งทดลองที่ 1 : อุณหภูมิทอดที่ 140 องศาเซลเซียส
 - สิ่งทดลองที่ 2 : อุณหภูมิทอดที่ 150 องศาเซลเซียส
 - สิ่งทดลองที่ 3 : อุณหภูมิทอดที่ 160 องศาเซลเซียส
- ผลตอบสนองหรือค่าสังเกต คือ ปริมาณไขมันในนั้กเกิดที่วิเคราะห์ได้

กรณีที่ 2 สิ่งทดลองหรือทริทเมนต์ประกอบด้วยหลายปัจจัย

นักวิจัยต้องการศึกษาปัจจัยในการทอด 2 ปัจจัยคือ อุณหภูมิในการทอด (150 และ 160 องศาเซลเซียส) และ เวลาในการทอด (3 และ 4 นาที) ว่ามีผลต่อการอมน้ำมันของนักเก็ตจากแป้งถั่วเหลืองหรือไม่ กรณีนี้เมื่อพิจารณาจะได้รายละเอียดของข้อมูล ดังนี้

- ปัจจัยที่ศึกษา มี 2 ปัจจัย คือ อุณหภูมิในการทอด และ เวลาในการทอด
- ระดับของแต่ละปัจจัย คือ
 - อุณหภูมิในการทอด มี 2 ระดับ (150 และ 160 องศาเซลเซียส)
 - เวลาในการทอด มี 2 ระดับ (3 และ 4 นาที)
- จำนวนสิ่งทดลองมี 4 สิ่งทดลอง (เท่ากับ ผลคูณของระดับแต่ละปัจจัย) ได้แก่
 - สิ่งทดลองที่ 1 : ใช้สภาวะทอดที่อุณหภูมิ 150 องศาเซลเซียส นาน 3 นาที
 - สิ่งทดลองที่ 2 : ใช้สภาวะทอดที่อุณหภูมิ 150 องศาเซลเซียส นาน 4 นาที
 - สิ่งทดลองที่ 3 : ใช้สภาวะทอดที่อุณหภูมิ 160 องศาเซลเซียส นาน 3 นาที
 - สิ่งทดลองที่ 4 : ใช้สภาวะทอดที่อุณหภูมิ 160 องศาเซลเซียส นาน 4 นาที
- ผลตอบสนองหรือค่าสังเกต คือ ปริมาณไขมันในนักเก็ตที่วิเคราะห์ได้

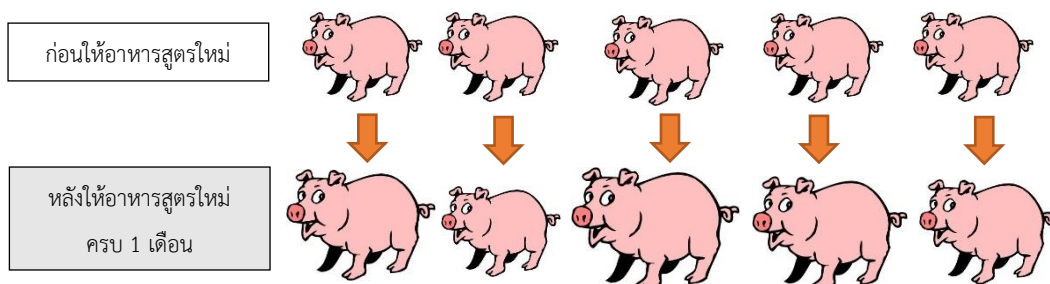
กรณีที่ 3 สิ่งทดลองหรือทริทเมนต์ที่ไม่ได้มาจากปัจจัยในระดับต่างๆ แต่มีลักษณะเฉพาะตัวของหน่วยทดลองแตกต่างกัน

นักวิจัยต้องการศึกษาว่าขนมขบเคี้ยวที่ใช้แป้งต่างชนิดกัน (แป้งข้าวเจ้า แป้งมันสำปะหลัง และแป้งสาลี) จะมีค่าความแข็งแตกต่างกันหรือไม่ กรณีนี้เมื่อพิจารณาจะได้รายละเอียดของข้อมูล ดังนี้

- ลักษณะประจำตัวของหน่วยทดลอง คือ ชนิดแป้งที่ใช้ทำขนมขบเคี้ยว
- จำนวนสิ่งทดลองมี 3 สิ่งทดลอง (เท่ากับลักษณะประจำตัวของหน่วยทดลอง) ได้แก่
 - สิ่งทดลองที่ 1 : แป้งข้าวเจ้า
 - สิ่งทดลองที่ 2 : แป้งมันสำปะหลัง
 - สิ่งทดลองที่ 3 : แป้งสาลี
- ผลตอบสนองหรือค่าสังเกต คือ ค่าความแข็งของขนมขบเคี้ยว

2.10 ความคลาดเคลื่อนในการทดลอง (Experimental error)

ความคลาดเคลื่อนในการทดลอง หมายถึง ความแปรผันระหว่างค่าสังเกตของหน่วยทดลองภายใต้สิ่งทดลองเดียวกัน ซึ่งความคลาดเคลื่อนนี้ไม่ได้เกิดจากอิทธิพลของสิ่งทดลองแต่เกิดจากปัจจัยภายนอกหรืออิทธิพลอื่นๆ ที่ไม่ทราบสาเหตุที่แน่ชัด เช่น โรงงานแห่งหนึ่งพัฒนาอาหารเลี้ยงสุกรสูตรใหม่และต้องการทราบว่าหากใช้อาหารสูตรใหม่นี้เลี้ยงสุกรเป็นเวลา 1 เดือน สุกรจะมีน้ำหนักเพิ่มขึ้นคิดเป็นร้อยละเท่าไร จึงทำการสุ่มสุกรสายพันธุ์เดียวกัน อายุและน้ำหนักเท่ากัน จำนวน 5 ตัว มาทดลอง ถ้าให้อาหารและน้ำเท่าๆ กัน สุกรทั้ง 5 ตัวนี้เมื่อครบ 1 เดือนควรมีน้ำหนักที่ใกล้เคียงกัน แต่ปรากฏว่าเมื่อครบกำหนด 1 เดือนน้ำหนักสุกรที่ได้มีความแตกต่างกัน (ดังภาพที่ 2.7) ความแปรผันนี้เรียกว่า **“ความคลาดเคลื่อนในการทดลอง”** ในการดำเนินงานวิจัยนั้นนักวิจัยต้องพยายามที่จะลดความคลาดเคลื่อนในการทดลองให้น้อยลงเท่าที่จะทำได้



ภาพที่ 2.7 การเปรียบเทียบความคลาดเคลื่อนในการทดลอง

2.11 การทำซ้ำ (Replication)

การทำซ้ำ คือ การใช้สิ่งทดลองเดิมทดลองซ้ำหรือกระทำซ้ำแก่หน่วยทดลองตั้งแต่ 2 หน่วยขึ้นไป ดังนั้นจำนวนครั้งของการทดลองในแต่ละสิ่งทดลองคือจำนวนหน่วยทดลองในแต่ละสิ่งทดลองนั่นเอง ทั้งนี้ในแต่ละสิ่งทดลองจะมีจำนวนซ้ำเท่ากันหรือไม่เท่ากันก็ได้ การทดลองในลักษณะเดียวกันหลายๆ ครั้ง หรือ การทดลองซ้ำหลายๆ ครั้งนี้จะช่วยให้ผู้ทดสอบสามารถประเมินค่าความคลาดเคลื่อนได้ และช่วยเพิ่มความไว (Sensitive) ในการวัดความแตกต่างระหว่างสิ่งทดลอง ทำให้การสรุปผลการทดลองสามารถทำได้กว้างขวางมากขึ้น ตัวอย่างกรณีการทำซ้ำ เช่น นักวิจัยต้องการให้ผู้ทดสอบประเมินความชอบเค้กสูตรใหม่ จึงให้ผู้ทดสอบจำนวน 50 คนชิมเค้กสูตรใหม่แล้วให้คะแนนความชอบ ในกรณีนี้สิ่งทดลองคือ ตัวอย่างเค้กสูตรใหม่ ผู้ทดสอบทั้ง 50 คน (เปรียบเสมือนหน่วยทดลอง 50 หน่วย) ได้รับตัวอย่างเค้กสูตรเดียวกัน จึงถือเป็นการทำซ้ำ 50 ครั้ง

จำนวนซ้ำของการทำซ้ำจะมากหรือน้อยขึ้นอยู่กับความเหมาะสมของหลายปัจจัย เช่น ความแปรผันของหน่วยทดลองถ้ามีสูงจะต้องใช้จำนวนซ้ำมาก, ความแตกต่างของสิ่งทดลองที่มีความแตกต่างกันมากจะใช้

จำนวนซ้ำน้อยกว่าสิ่งทดลองที่มีความแตกต่างกันน้อย, การทดลองที่ใช้การทดสอบแบบทางเดียวจะใช้จำนวนซ้ำน้อยกว่าการทดสอบแบบสองทาง และ ปัจจัยด้านงบประมาณ แรงงาน และ ระยะเวลา

2.12 การสุ่ม (Randomization)

การสุ่ม คือ วิธีการจัดสิ่งทดลองให้มีโอกาสเท่าๆ กันให้แก่หน่วยทดลองด้วยวิธีการสุ่ม การสุ่มจะทำให้แต่ละสิ่งทดลองถูกจัดลงในหน่วยทดลองอย่างไม่มีอคติ และช่วยให้แน่ใจว่าการประมาณค่าความคลาดเคลื่อนในการทดลอง การประมาณค่าเฉลี่ยของสิ่งทดลอง หรือการหาความแตกต่างระหว่างสิ่งทดลองเป็นไปอย่างสมเหตุสมผล ยกตัวอย่างเช่น ในการทดสอบทางด้านประสาทสัมผัส ลำดับการเสิร์ฟตัวอย่างมีผลต่อความรู้สึกของผู้ทดสอบ ดังนั้นหากมีตัวอย่าง 2 ตัวอย่าง (ตัวอย่าง A และตัวอย่าง B) นักวิจัยต้องทำการสุ่มลำดับการเสิร์ฟตัวอย่าง A และตัวอย่าง B ให้มีโอกาสดำเนินการอยู่ตำแหน่งที่ 1 และ 2 ดังแสดงในตารางที่ 2.4 เมื่อพิจารณาจะพบว่าตัวอย่าง A และตัวอย่าง B ถูกเสิร์ฟในลำดับที่ 1 จำนวน 15 ครั้ง เท่ากัน

ตารางที่ 2.4 ลำดับการเสิร์ฟตัวอย่างให้ผู้ทดสอบประเมินความชอบ

ผู้ทดสอบคนที่	ลำดับเสิร์ฟที่ 1	ลำดับเสิร์ฟที่ 2	ผู้ทดสอบคนที่	ลำดับเสิร์ฟที่ 1	ลำดับเสิร์ฟที่ 2
1	ตัวอย่าง A	ตัวอย่าง B	16	ตัวอย่าง A	ตัวอย่าง B
2	ตัวอย่าง B	ตัวอย่าง A	17	ตัวอย่าง A	ตัวอย่าง B
3	ตัวอย่าง A	ตัวอย่าง B	18	ตัวอย่าง A	ตัวอย่าง B
4	ตัวอย่าง B	ตัวอย่าง A	19	ตัวอย่าง B	ตัวอย่าง A
5	ตัวอย่าง A	ตัวอย่าง B	20	ตัวอย่าง A	ตัวอย่าง B
6	ตัวอย่าง B	ตัวอย่าง A	21	ตัวอย่าง B	ตัวอย่าง A
7	ตัวอย่าง B	ตัวอย่าง A	22	ตัวอย่าง A	ตัวอย่าง B
8	ตัวอย่าง B	ตัวอย่าง A	23	ตัวอย่าง B	ตัวอย่าง A
9	ตัวอย่าง B	ตัวอย่าง A	24	ตัวอย่าง B	ตัวอย่าง A
10	ตัวอย่าง A	ตัวอย่าง B	25	ตัวอย่าง A	ตัวอย่าง B
11	ตัวอย่าง A	ตัวอย่าง B	26	ตัวอย่าง B	ตัวอย่าง A
12	ตัวอย่าง A	ตัวอย่าง B	27	ตัวอย่าง B	ตัวอย่าง A
13	ตัวอย่าง A	ตัวอย่าง B	28	ตัวอย่าง B	ตัวอย่าง A
14	ตัวอย่าง B	ตัวอย่าง A	29	ตัวอย่าง A	ตัวอย่าง B
15	ตัวอย่าง A	ตัวอย่าง B	30	ตัวอย่าง B	ตัวอย่าง A

2.13 ความสม่ำเสมอของหน่วยทดลอง (Uniformity)

ความสม่ำเสมอของหน่วยทดลอง คือ การที่หน่วยทดลองมีลักษณะเฉพาะตัวเหมือนกัน หน่วยทดลองที่เหมือนกันเมื่อได้รับสิ่งทดลองเดียวกันจะให้ผลที่คล้ายคลึงกัน การใช้หน่วยทดลองที่เหมือนกันจะช่วยลดความคลาดเคลื่อนของการทดลอง

- กรณีหน่วยทดลองเป็นคน เช่น ต้องการประเมินคุณภาพทางด้านประสาทสัมผัสควรใช้ผู้ทดสอบที่มีช่วงอายุเดียวกัน ในการทดสอบความชอบควรใช้ผู้ทดสอบทั่วไปที่ไม่ได้รับการฝึกฝน (Untrained panelists) และไม่ควรใช้ผู้ที่คุ้นเคยกับผลิตภัณฑ์เป็นอย่างดี เช่น พนักงานควบคุมคุณภาพ
- กรณีหน่วยทดลองเป็นสัตว์ เช่น ต้องการทดสอบสูตรอาหารควรเลือกใช้สัตว์ชนิดเดียว สายพันธุ์เดียวกัน เพศเดียวกัน อายุ และ น้ำหนักใกล้เคียงกัน
- กรณีหน่วยทดลองเป็นพืช เช่น ต้องการทดสอบสูตรปุ๋ยควรเลือกพืชชนิดเดียวกัน อายุใกล้เคียงกัน ส่วนสูงและขนาดลำต้นใกล้เคียงกัน ปลูกในพื้นที่เพาะปลูกเดียวกัน ได้รับแสงและน้ำที่เท่ากัน

2.14 การบล็อก (Blocking)

การบล็อก คือ การจัดหน่วยทดลองที่มีความคล้ายคลึงกันหรือเหมือนกันให้อยู่ในบล็อกเดียวกันหรือกลุ่มเดียวกัน เพื่อให้มีความแปรปรวนภายในบล็อกเดียวน้อยที่สุด แต่หน่วยทดลองที่อยู่ต่างบล็อกกันจะมีความแตกต่างกันมากที่สุด เพื่อให้มีความแปรปรวนระหว่างบล็อกมากที่สุด เช่น นักวิจัยต้องการศึกษาอุณหภูมิในการทอดว่ามีผลอย่างไรต่อการอมน้ำมันของเฟืองแผ่นทอด เมื่อซื้อเฟืองมา 10 กิโลกรัมพบว่าสามารถจำแนกเฟืองออกตามขนาดได้ 3 ขนาด ได้แก่ เล็ก กลาง และ ใหญ่ จึงทำการบล็อกเฟือง 3 ขนาด ได้แก่ บล็อกที่ 1 คือ เฟืองขนาดเล็ก บล็อกที่ 2 คือ เฟืองขนาดกลาง และ บล็อกที่ 3 คือเฟืองขนาดใหญ่

ในการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (Randomized Completely Block Design หรือ RCBD) บล็อก คือ การทำซ้ำของการทดลอง เช่นกรณีตัวอย่างการศึกษาดูอุณหภูมิในการทอดว่ามีผลอย่างไรต่อการอมน้ำมันของเฟืองแผ่นทอด หากสนใจอุณหภูมิในการทอด 2 ระดับ คือ 150 และ 160 องศาเซลเซียส จะกำหนดให้แต่ละอุณหภูมิการทอดทำการทดลองซ้ำ 3 ซ้ำการทดลอง กล่าวคือ ซ้ำที่ 1 ใช้เฟืองขนาดเล็ก (บล็อกที่ 1) ซ้ำที่ 2 ใช้เฟืองขนาดกลาง (บล็อกที่ 2) และซ้ำที่ 3 ใช้เฟืองขนาดใหญ่ (บล็อกที่ 3)



บทที่ 3

ความรู้พื้นฐานเกี่ยวกับการใช้ โปรแกรมสำเร็จรูป SPSS for Windows

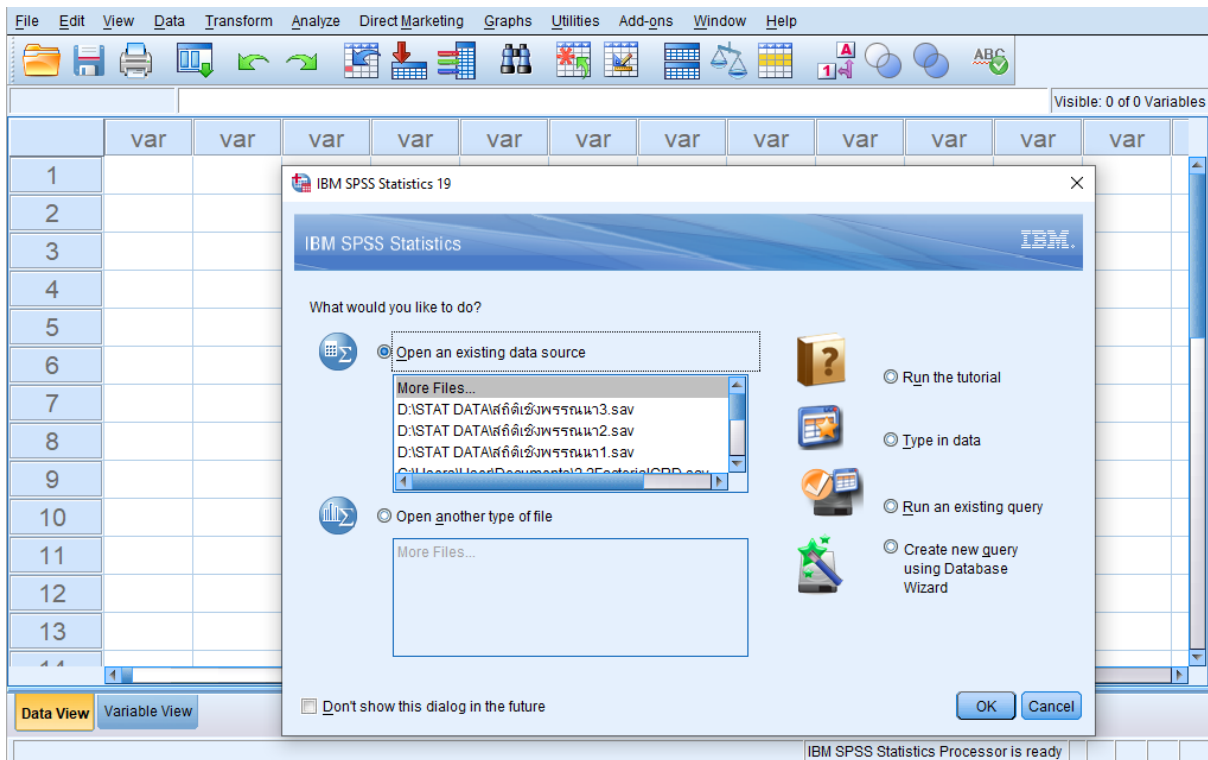
SPSS for Windows (Statistical Package for Social Science) เป็นโปรแกรมสำเร็จรูปทางสถิติที่ใช้ในการวิเคราะห์ข้อมูลทั้งสถิติพื้นฐานและสถิติขั้นสูงที่มีประสิทธิภาพ เป็นโปรแกรมทางสถิติที่มีการจัดการข้อมูลอย่างมีระบบ มีการแสดงค่าสถิติต่างๆ พร้อมกราฟและตาราง อย่างไรก็ตามผู้ใช้งานควรมีความรู้ความเข้าใจในการจัดการข้อมูล การเตรียมข้อมูล การป้อนข้อมูลให้เหมาะสมกับวิธีการวิเคราะห์ทางสถิติที่เลือกใช้ ตลอดจนการแปลความหมายค่าสถิติจากการวิเคราะห์เพื่อสรุปผลการวิเคราะห์ทางสถิติได้อย่างถูกต้อง

ในงานวิจัยด้านพัฒนาผลิตภัณฑ์สามารถใช้โปรแกรม SPSS ได้ในหลายขั้นตอนของกระบวนการพัฒนาผลิตภัณฑ์ เช่น การสร้างแฟ้มข้อมูลจากแบบสอบถามที่ผู้วิจัยใช้ในการสำรวจพฤติกรรมและความต้องการของผู้บริโภค ตลอดจนการวิเคราะห์หาความสัมพันธ์ระหว่างกลุ่มผู้บริโภคและพฤติกรรมผู้บริโภค หรือใช้เพื่อวิเคราะห์สรุปผลการวิจัยตามวัตถุประสงค์ เช่น การคัดเลือกสูตรที่เหมาะสม การเปรียบเทียบความแตกต่างระหว่างวัตถุดิบเดิมและวัตถุดิบชนิดใหม่ที่มีผลต่อคุณภาพของผลิตภัณฑ์ การหาสภาวะที่เหมาะสมของกระบวนการผลิต การสร้างสมการหาความสัมพันธ์ระหว่างตัวแปรและการพยากรณ์ค่าในอนาคต

ในหนังสือบทนี้ได้รวบรวมความรู้พื้นฐานเกี่ยวกับการใช้โปรแกรม SPSS for Windows version 19 ได้แก่ การเริ่มต้นใช้งานโปรแกรม ข้อกำหนดทั่วไปของโปรแกรม การป้อนข้อมูล และแนะนำเมนูสำหรับการวิเคราะห์ทางสถิติ ทั้งนี้คำสั่งของโปรแกรม SPSS สำหรับการวิเคราะห์ค่าทางสถิติแต่ละวิธีจะแยกอยู่ตามบทต่างๆ อย่างไรก็ตามหากนักวิจัยใช้โปรแกรม SPSS version อื่นก็สามารถใช้หลักการเดียวกันในการวิเคราะห์ทางสถิติได้เช่นกัน

3.1 การเริ่มต้นใช้งานโปรแกรม SPSS for Windows

กดดับเบิลคลิก (Double click) เปิดโปรแกรม **IBM SPSS Statistics 19** จะเข้าสู่โปรแกรมดังแสดงในภาพที่ 3.1



ภาพที่ 3.1 การเริ่มต้นเข้าสู่โปรแกรม SPSS

จากภาพที่ 3.1 ให้สังเกตคำสั่งในหน้าจอ ผู้ใช้งานสามารถเริ่มต้นใช้งานได้หลายแบบ ดังนี้

แบบที่ 1 กรณีผู้วิจัยต้องการเลือกใช้แฟ้มข้อมูลงานเดิมที่เคยป้อนข้อมูลไว้ในโปรแกรม SPSS (ไฟล์ที่มีนามสกุล .sav) ทำได้โดยกดคลิกที่ **Open an existing data source** แล้วทำการหาแฟ้มงานเดิมที่ต้องการใช้งาน

แบบที่ 2 กรณีผู้วิจัยต้องการเลือกใช้แฟ้มข้อมูลงานเดิมที่เคยป้อนข้อมูลในโปรแกรมอื่นที่เป็นไฟล์ประเภท Spread sheet ทำได้โดยกดคลิกที่ **Open another type of file** อ่านรายละเอียดในข้อ 3.3

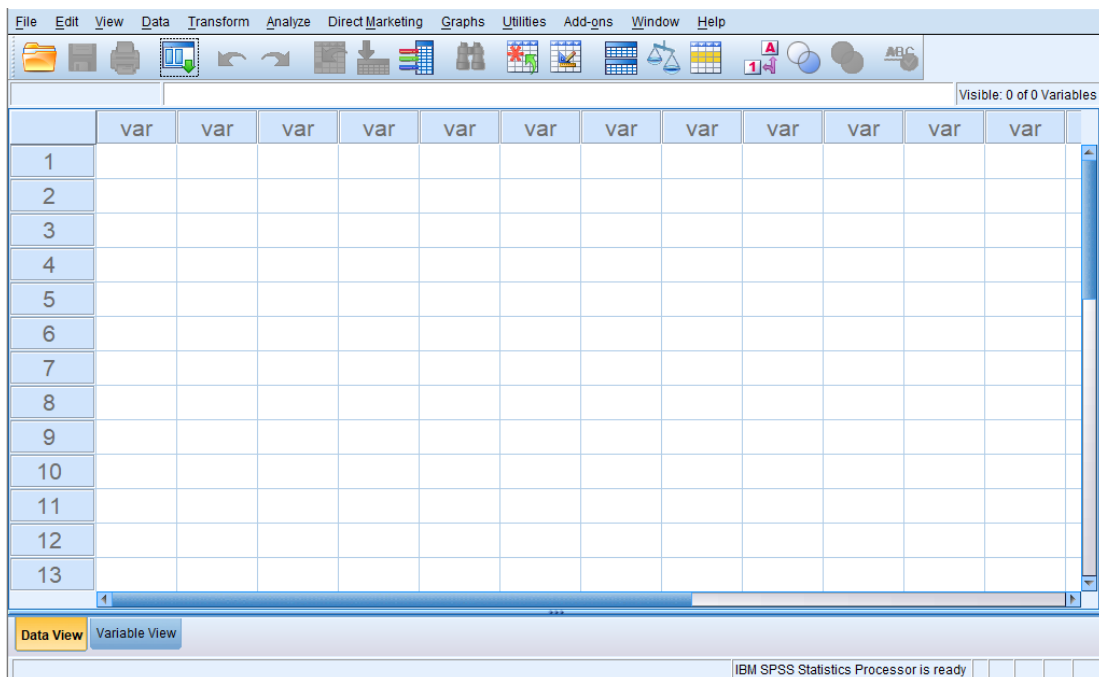
แบบที่ 3 กรณีผู้วิจัยต้องการพิมพ์ข้อมูลงานใหม่ ให้เลือกคลิก **Type in data** แล้วกดคลิก OK

3.2 ข้อกำหนดทั่วไปของโปรแกรม SPSS for Windows

โปรแกรม SPSS แบ่งส่วนการใช้งานออกเป็น 2 หน้าต่าง (Sheet) สังเกตที่ด้านล่างของหน้าจอ (ภาพที่ 3.2) จะปรากฏชื่อของทั้งสองหน้าต่าง คือ **Data View** และ **Variable View** โดยแต่ละหน้าต่างจะทำหน้าที่ต่างกัน ดังนี้

3.2.1 หน้าต่าง Data View สำหรับใช้ในการป้อนข้อมูล ดังแสดงในภาพที่ 3.2

ด้านซ้ายของตารางแสดงตัวเลข 1 2 3 4 ... แสดงถึง case หรือข้อมูลตัวที่ 1 2 3 4 ... ตามลำดับ และ ด้านบนของตารางในแนวคอลัมน์จะปรากฏชื่อตัวแปรเป็น var เมื่อทำการป้อนข้อมูลลงในตาราง ชื่อตัวแปร var จะเปลี่ยนเป็น VAR00001 VAR00002 VAR00003 ... ซึ่งผู้ใช้งานสามารถกำหนดชื่อให้ตัวแปรใหม่ได้ โดยทำการคลิกที่หน้าต่าง Variable View ที่หน้าจอด้านล่าง (ซึ่งจะกล่าวถึงต่อไปในข้อ 3.2.2)



ภาพที่ 3.2 หน้าต่าง Data View ในโปรแกรม SPSS

การเพิ่มจำนวนตัวแปร var หรือ จำนวน case สามารถทำได้ 2 วิธี คือ

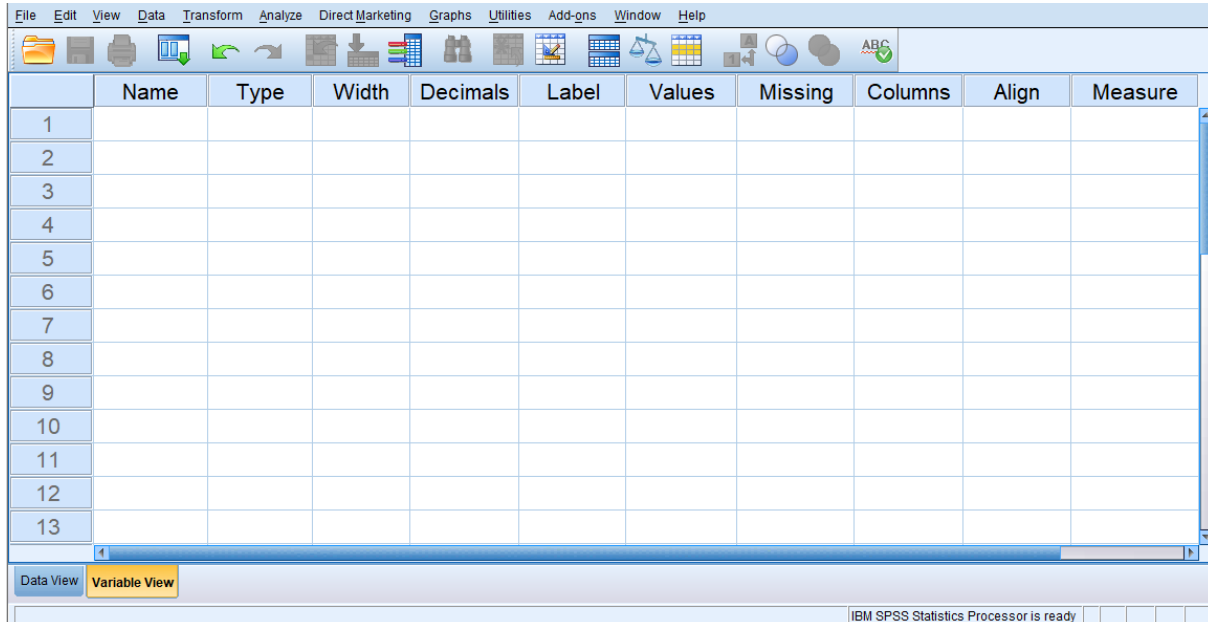
วิธีที่ 1 ใช้คำสั่ง **Edit** ที่เมนูคำสั่งด้านบน

(1) กรณีต้องการเพิ่มจำนวนตัวแปร ให้กดคลิกที่ตัวแปร var ด้านบนคอลัมน์ให้เป็นแถบสี จากนั้นใช้คำสั่ง **Edit** ที่เมนูคำสั่งด้านบน คลิกเลือก **Insert Variable** จะปรากฏตัวแปร var ใหม่เพิ่มขึ้นมาหน้าตัวแปร var เดิม

(2) กรณีต้องการเพิ่มจำนวน case ให้กดคลิกที่ตัวเลข case ซ้ายมือให้เป็นแถบสี จากนั้นใช้คำสั่ง **Edit** ที่เมนูคำสั่งด้านบน คลิกเลือก **Insert Cases** จะปรากฏ case ใหม่เพิ่มขึ้นมา

วิธีที่ 2 ถ้าต้องการเพิ่มจำนวนตัวแปร var ให้กดคลิกที่คอลัมน์ var ให้เป็นแถบสี และคลิกเมาส์ปุ่มขวาเลือก **Insert Variable** และถ้าต้องการเพิ่มจำนวน case ให้กดคลิกที่ตัวเลขในแถวบน ให้เป็นแถบสี และคลิกเมาส์ปุ่มขวาเลือก **Insert Cases** ถ้าต้องการเพิ่มจำนวน case

3.2.2 หน้าต่าง Variable View สำหรับใช้ในการกำหนดรายละเอียดของตัวแปร เช่น ชื่อของตัวแปร (Name) ชนิดของตัวแปร (Type) ขนาดความกว้างของคอลัมน์ (Width) จำนวนทศนิยม (Decimals) ดังแสดงในภาพที่ 3.3



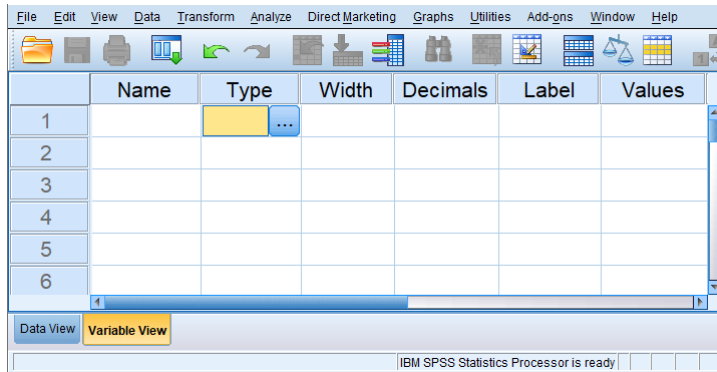
ภาพที่ 3.3 หน้าต่าง Variable View ในโปรแกรม SPSS

ด้านซ้ายของหน้าต่าง Variable View แสดงตัวเลข 1 2 3 ... ซึ่งหมายถึง ตัวแปร var ตัวที่ 1 2 3 ... ตามลำดับ ด้านบนของหน้าต่าง Variable View ในแต่ละคอลัมน์จะใช้ในการกำหนดรายละเอียดของตัวแปร ดังนี้

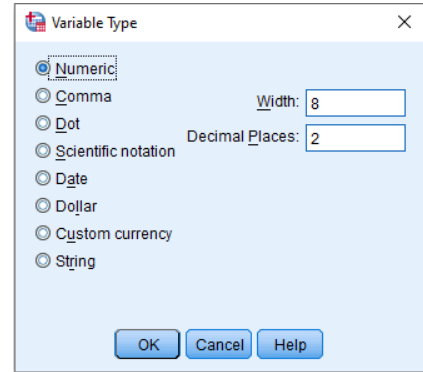
(1) Name ใช้ในการกำหนดชื่อ ตัวแปร var มีข้อกำหนด ดังนี้

- มีได้ไม่เกิน 64 ตัวอักษร (รวมตัวอักษรและตัวเลข)
- ห้ามเว้นช่องว่างหรือเว้นวรรค
- ในแฟ้มข้อมูลเดียวกันแต่ละตัวแปรห้ามใช้ชื่อซ้ำ
- ห้ามมีสัญลักษณ์อื่นๆ เช่น () + - × / *
- ชื่อตัวแปรต้องไม่จบด้วยจุด ตัวอักษรตัวเล็กหรือตัวใหญ่ถือเป็นตัวแปรเดียวกัน เช่น TEMP Temp temp ถือเป็นตัวแปรเดียวกัน
- ชื่อตัวแปรต้องเริ่มต้นด้วยตัวอักษรเท่านั้น
- ห้ามใช้ ALL, AND, BY, EQ, GE, GT, LE, LT, NE, NOT, OR, TO และ WITH ในการตั้งชื่อตัวแปร

(2) **Type** ใช้ในการกำหนดชนิดของตัวแปร เมื่อกดที่ Cell ในคอลัมน์ของ Type (ภาพที่ 3.4) จะปรากฏกล่องเมนู Variable Type ซึ่งมีชนิดของตัวแปรให้เลือกดังภาพที่ 3.5 เช่น Numeric (ตัวแปรชนิดตัวเลข), String (ตัวแปรที่มีค่าเป็นอักษร ตัวเลข หรือเครื่องหมายต่างๆ), Date (ตัวแปรชนิดวันที่) โดยปกติโปรแกรมจะกำหนดไว้เป็นข้อมูล Numeric ทั้งนี้ตัวแปรหนึ่งๆ จะต้องเป็นชนิดใดชนิดหนึ่งเท่านั้น

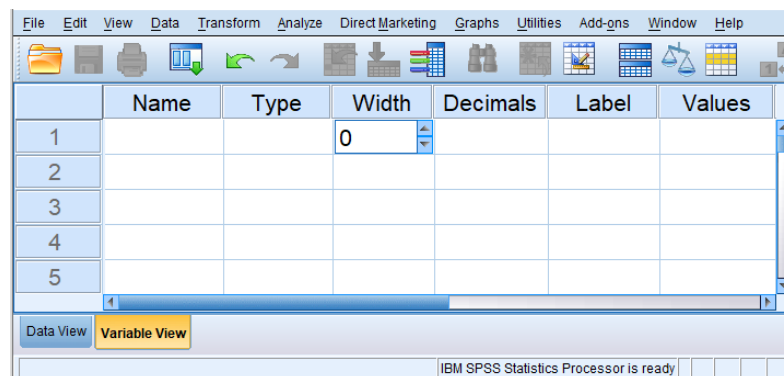


ภาพที่ 3.4 การกำหนดชนิดตัวแปร (Type)



ภาพที่ 3.5 เมนู Variable Type

(3) **Width** ใช้ในการกำหนดความกว้างหรือจำนวนหลักของค่าตัวแปร โปรแกรมจะกำหนดค่า Width อัตโนมัติไว้เท่ากับ 8 ซึ่งนักวิจัยสามารถเปลี่ยนได้ตามความเหมาะสมของข้อมูลโดยกดที่ Cell ในคอลัมน์ของ Width แล้วคลิกลูกศรเพิ่มหรือลดจำนวนหลักของข้อมูล (ภาพที่ 3.6)



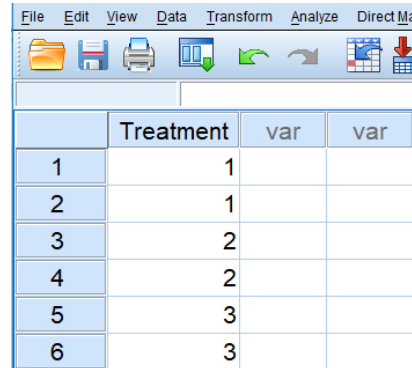
ภาพที่ 3.6 การกำหนดความกว้างของค่าตัวแปรใน Width

(4) **Decimals** ใช้ในการกำหนดจำนวนจุดทศนิยมของตัวแปร บางครั้งนักวิจัยต้องการป้อนข้อมูลที่เป็นเลขจำนวนเต็ม เช่น สิ่งทดลองที่ 1 2 3 แต่โดยปกติโปรแกรมจะตั้งค่าไว้อัตโนมัติที่ทศนิยม 2 ตำแหน่ง เมื่อเวลาป้อนข้อมูลที่หน้าต่าง Data View จะปรากฏเป็น 1.00 2.00 3.00 ดังภาพที่ 3.7 เนื่องจากตัวเลขที่ระบุสิ่งทดลองที่ 1 2 3 เป็นข้อมูลชนิดนามกำหนด (Nominal) ดังนั้นนักวิจัยจึงไม่ควรแสดงข้อมูลที่มีจุดทศนิยม โดยสามารถตั้งค่าจำนวนจุดทศนิยมให้เป็น 0 ได้โดยกดที่ Cell ในคอลัมน์ของ Decimals (หน้าต่าง

Variable View) แล้วคลิกลูกศรลดจำนวนให้เท่ากับ 0 (ภาพที่ 3.9) ซึ่งจะทำให้ได้หน้าต่าง Data View ปรากฏเป็นข้อมูลเลขจำนวนเต็ม ดังภาพที่ 3.8



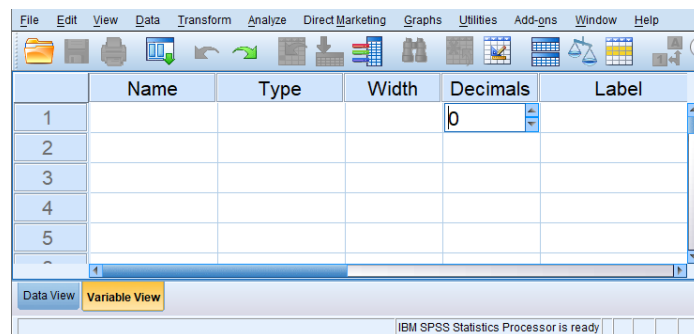
	Treatment	var	var
1	1.00		
2	1.00		
3	2.00		
4	2.00		
5	3.00		
6	3.00		



	Treatment	var	var
1	1		
2	1		
3	2		
4	2		
5	3		
6	3		

ภาพที่ 3.7 หน้าจอเมื่อโปรแกรมกำหนดทศนิยมอัตโนมัติ

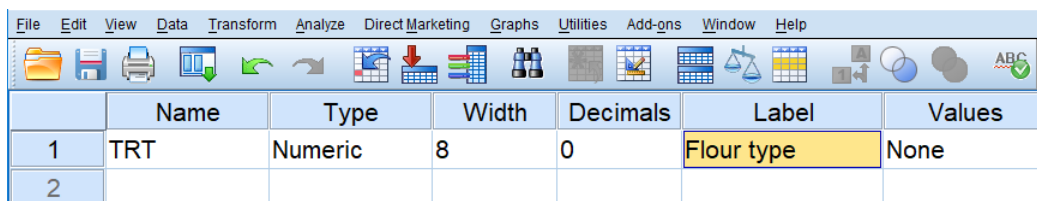
ภาพที่ 3.8 หน้าจอเมื่อกำหนดทศนิยมให้เป็นศูนย์



	Name	Type	Width	Decimals	Label
1				0	
2					
3					
4					
5					

ภาพที่ 3.9 การกำหนดจำนวนทศนิยมให้ค่าตัวแปร Var ใน Decimals

(5) **Label** ใช้ในการอธิบายความหมายของตัวแปร (var) เช่น ผู้ใช้งานกำหนดชื่อตัวแปร var เป็น TRT และต้องการอธิบายความหมายที่แท้จริงของ TRT ว่าหมายถึง ชนิดของแป้ง (Flour type) นักวิจัยสามารถระบุความหมายของตัวแปรได้โดยกดที่ Cell ในคอลัมน์ของ Label แล้วพิมพ์ความหมาย (ดังภาพที่ 3.10) ข้อดีของการอธิบายความหมายของตัวแปร var ใน Label คือ ความหมายของตัวแปรนี้จะแสดงในผลลัพธ์ (Output) ที่ได้จากการวิเคราะห์ทางสถิติ ทำให้ลดความผิดพลาดในการรายงานผลการวิเคราะห์ได้



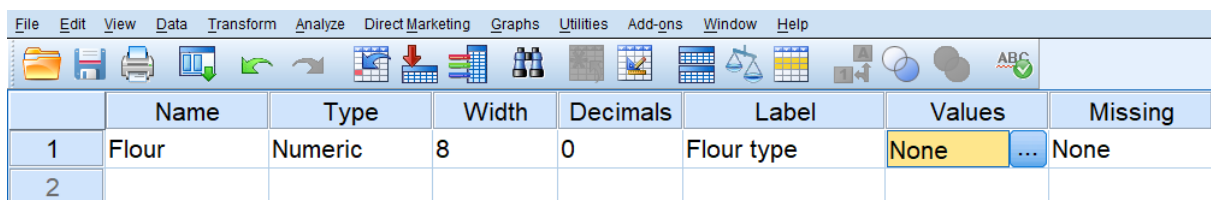
	Name	Type	Width	Decimals	Label	Values
1	TRT	Numeric	8	0	Flour type	None
2						

ภาพที่ 3.10 การกำหนดคำอธิบายความหมายของตัวแปร Var ใน Label

(6) **Values** ใช้ในการกำหนดค่าของตัวแปร มักใช้ในกรณีที่แปลงข้อมูลเชิงกลุ่มให้เป็นตัวเลข ยกตัวอย่างเช่น นักวิจัยต้องการศึกษาผลของชนิดแป้งที่มีต่อความนุ่มของขนมปัง จึงทำการผลิตขนมปังจากแป้ง 3 ชนิด ได้แก่ ขนมปังจากแป้งสาลี, ขนมปังจากแป้งข้าว และขนมปังจากแป้งมันสำปะหลัง ในกรณีนี้ นักวิจัยจะมีสิ่งทดลองจำนวน 3 สิ่งทดลอง เมื่อต้องการป้อนข้อมูลสิ่งทดลองจึงกำหนดชื่อตัวแปร var เป็น Flour โดยกำหนดรหัสตัวเลขให้แต่ละสิ่งทดลอง ได้แก่ 1 = แป้งสาลี, 2 = แป้งข้าว และ 3 = แป้งมันสำปะหลัง

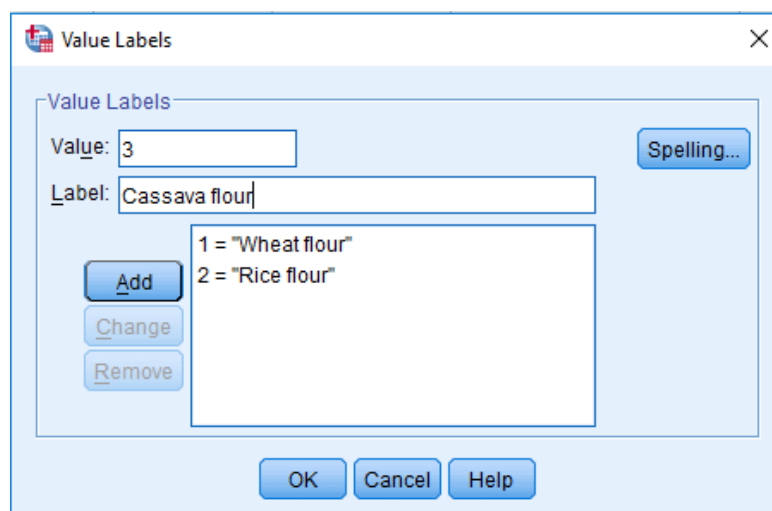
การกำหนดรหัสหรือค่าให้ตัวแปร ทำได้โดยกดที่ Cell ในคอลัมน์ของ Values แล้วกดคลิกที่มุมขวา (ภาพที่ 3.11) จะปรากฏกล่อง Value Labels ดังภาพที่ 3.12 ให้พิมพ์ค่า Value = 1 และ Label = Wheat flour และกด add จะปรากฏ 1 = Wheat flour ในกล่องด้านล่าง ต่อจากนั้นให้กำหนดค่าให้แป้งข้าว (2 = Rice flour) และ แป้งมันสำปะหลัง (3 = Cassava flour) เมื่อกำหนดค่า Values ทั้งหมดเสร็จแล้วให้กด OK

ข้อดีของการกำหนดค่าให้ตัวแปร var ใน Values คือ ค่าของตัวแปรนี้จะแสดงในผลลัพธ์ (Output) โดยที่ผู้วิจัยไม่ต้องจำว่า 1, 2, 3 ... คืออะไร ทำให้ลดความผิดพลาดในการสรุปผลการวิเคราะห์ได้อีกทางหนึ่ง



	Name	Type	Width	Decimals	Label	Values	Missing
1	Flour	Numeric	8	0	Flour type	None	None
2							

ภาพที่ 3.11 การกำหนดค่าให้กับตัวแปรใน Values



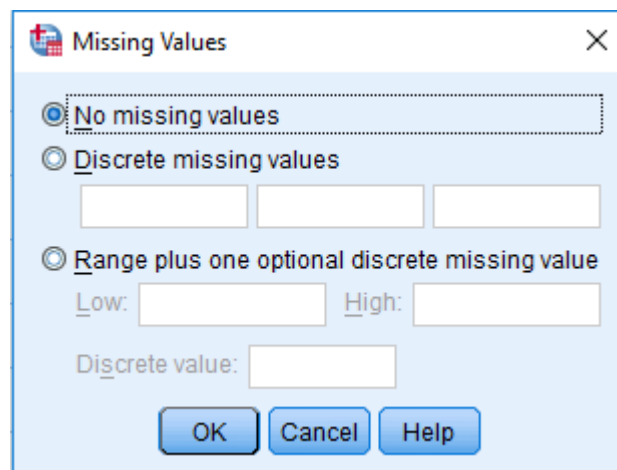
ภาพที่ 3.12 ตัวอย่างการกำหนดค่าให้กับตัวแปรใน Values

(7) **Missing** ใช้ในการกำหนดรหัสสำหรับค่าสูญหาย (Missing data) หรือ กรณีที่ผู้ตอบแบบสอบถามไม่ตอบ หรือผู้ใช้งานพิมพ์ข้อมูลไม่ครบ สามารถกำหนดได้ 3 แบบ ดังภาพที่ 3.13

แบบที่ 1 No missing values กรณีผู้วิจัยไม่พิมพ์ข้อมูล โปรแกรมจะกำหนดให้ค่าเป็นจุด (.)

แบบที่ 2 Discrete missing values กรณีผู้วิจัยเป็นผู้กำหนดรหัสของ Missing เอง

แบบที่ 3 Range plus one optional discrete missing value กรณีที่ผู้วิจัยกำหนดให้ผู้ตอบแบบสอบถามข้ามหรือไม่ต้องตอบคำถามบางข้อ ให้กำหนดรหัสของคำถามที่ต้องข้ามนั้นไว้อีกหนึ่งรหัสหนึ่ง



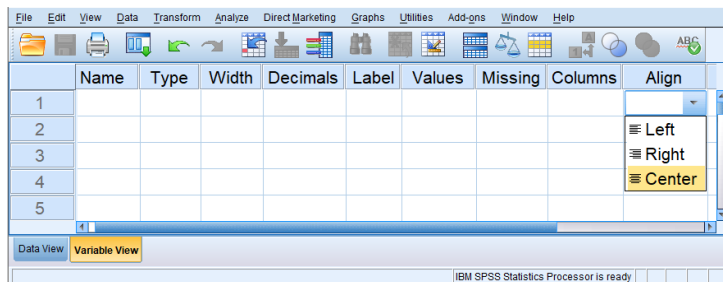
ภาพที่ 3.13 เมนูการกำหนดค่า Missing Values

(8) **Columns** ใช้ในการกำหนดความกว้างของ Columns เฉพาะในหน้าต่าง Data View ปกติ โปรแกรมจะกำหนดค่า Columns ไว้อัตโนมัติเท่ากับ 8 เท่ากับความกว้างของ Width (ดังภาพที่ 3.14) ทั้งนี้ นักวิจัยสามารถเปลี่ยนได้ตามความเหมาะสมของข้อมูลโดยกดที่ Cell ในคอลัมน์ของ Columns แล้วคลิก ลูกศรเพิ่มหรือลดจำนวนความกว้างของข้อมูล หากนักวิจัยกำหนดความกว้างของ Columns น้อยกว่าความกว้างของ Width ในหน้าต่าง Data View จะแสดงข้อมูลไม่ครบ อย่างไรก็ตามไม่มีผลกระทบต่อการวิเคราะห์ทางสถิติ

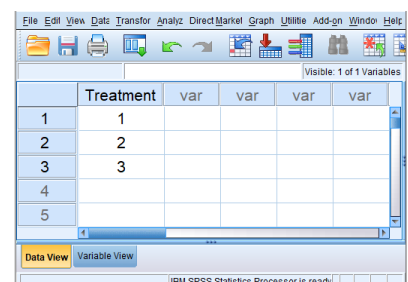
File Edit View Data Transform Analyze Direct Marketing Graphs Utilities Add-ons Window Help								
	Name	Type	Width	Decimals	Label	Values	Missing	Columns
1	Flour	Numeric	8	0	Flour type	{1, Wheat fl...	None	8
2								

ภาพที่ 3.14 ค่า Width และค่า Columns ที่โปรแกรมกำหนดไว้อัตโนมัติ

(9) **Align** ใช้ในการกำหนดตำแหน่งของข้อมูลใน Cell ของแฟ้มข้อมูลในหน้าต่าง Data View นักวิจัยสามารถจัดให้ชิดซ้าย (Left), ขวา (Right) หรืออยู่กลาง (Center) โดยกดที่ Cell ในคอลัมน์ของ Align ในหน้าต่าง Variable View จะปรากฏตัวเลือกดังภาพที่ 3.15a การกำหนดตำแหน่งของข้อมูลจะมีผลเฉพาะรูปแบบการจัดวางของข้อมูลในหน้าต่าง Data View เท่านั้น (ดังภาพที่ 3.15b) ไม่มีผลกระทบต่อการวิเคราะห์ทางสถิติ



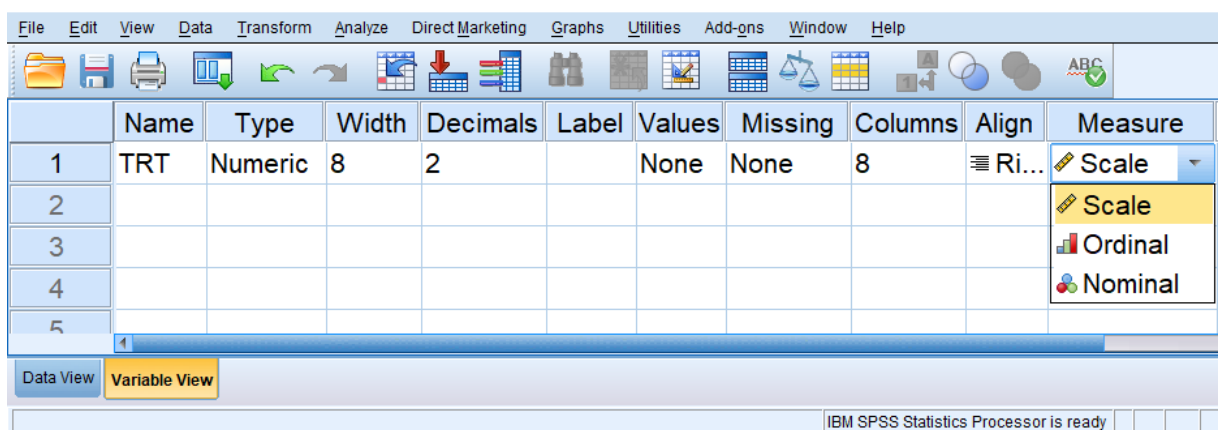
(a)



(b)

ภาพที่ 3.15 การกำหนดตำแหน่งของข้อมูลใน Align (a) และ
ตัวอย่างเมื่อกำหนด Center ให้ข้อมูลอยู่ตำแหน่งกลางของ Cell (b)

(10) **Measure** ใช้ในการกำหนดสเกลของข้อมูล ในโปรแกรม SPSS แบ่งเป็น 3 สเกล คือ Nominal, Ordinal และ Scale (หมายถึง Interval และ Ratio) ดังภาพที่ 3.16 (อ่านเรื่องประเภทของข้อมูลและสเกลในการวัดเพิ่มเติมได้ในบทที่ 2)

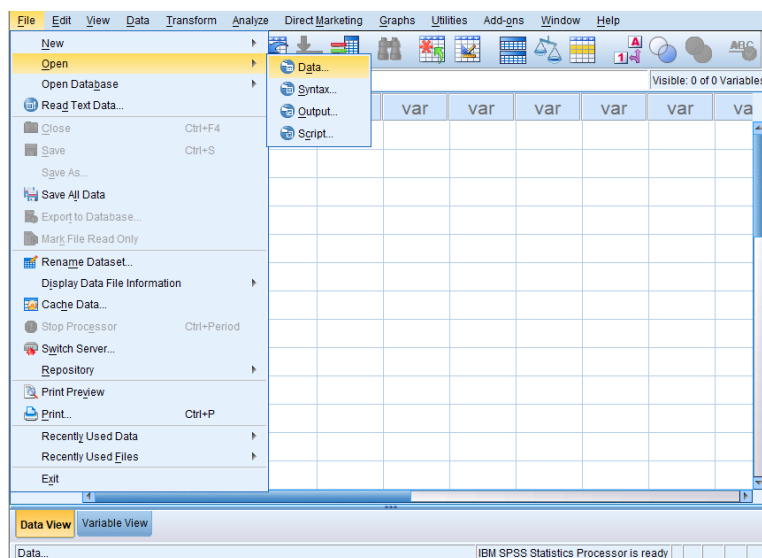


ภาพที่ 3.16 การกำหนดสเกลของข้อมูลใน Measure

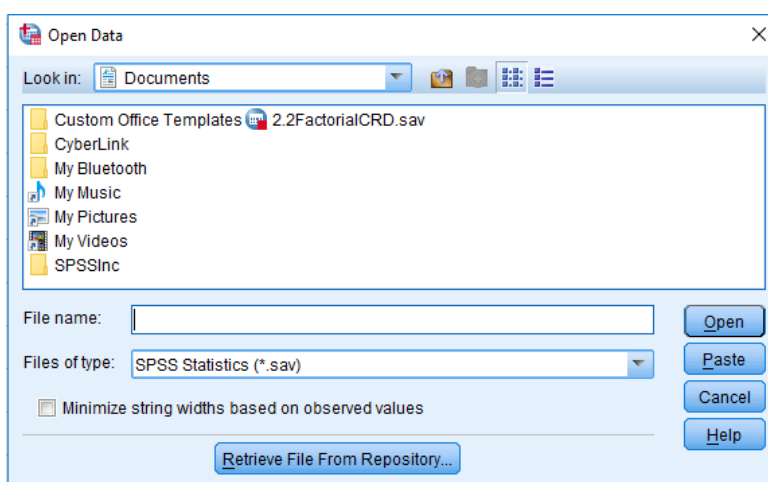
3.3 การป้อนข้อมูลหรือนำข้อมูลเข้าสู่ Work sheet

การป้อนข้อมูลหรือนำข้อมูลเข้าสู่ Work sheet ในโปรแกรม SPSS สามารถทำได้หลายวิธี แต่ 2 วิธีที่นิยมได้แก่ การเปิดจากไฟล์ประเภทอื่น และ การป้อนข้อมูลโดยตรง

3.3.1 การเปิดจากไฟล์ประเภทอื่น โปรแกรม SPSS สามารถเปิดไฟล์ได้จากโปรแกรมหลายประเภท ได้แก่ โปรแกรม Excel, Systat, Portable, Lotus, Sulk, dBase, SAS, Stata และ Text โดยนักวิจัยสามารถเลือกเปิดไฟล์ได้จากเมนูด้านบน **File ➤ Open ➤ Data** (ดังภาพที่ 3.17) จะปรากฏกล่อง Open Data ดังภาพที่ 3.18 ให้กดคลิกที่ Files of type เพื่อเลือกชนิดของไฟล์ และกดที่ Look in เพื่อเลือก Drive และ Folder ที่มีไฟล์ที่ต้องการจะเปิด



ภาพที่ 3.17 การเปิดไฟล์งานจากเมนู File / Open / Data

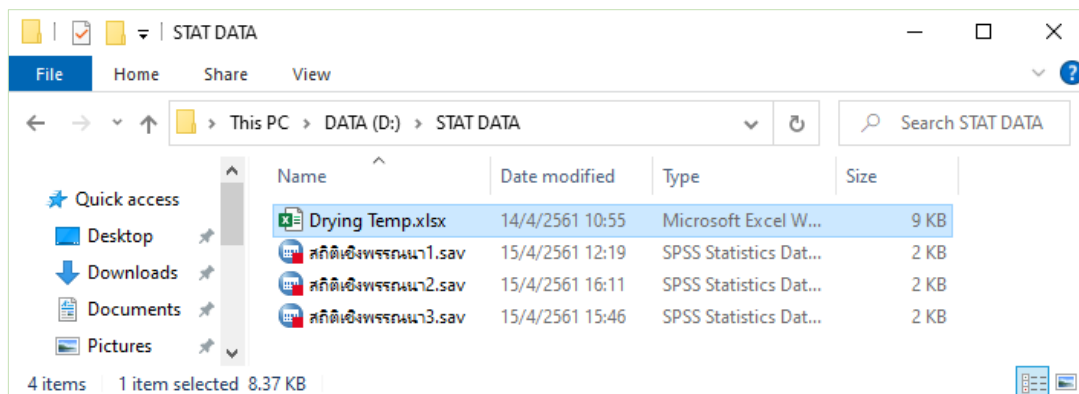


ภาพที่ 3.18 กล่อง Open Data จากเมนูคำสั่ง File / Open / Data

ในกรณีที่เลือกชนิดของไฟล์จากโปรแกรม Excel (*.xls, *.xlsx, *.xlsm) ซึ่งเป็นไฟล์ประเภท Spread sheet และเป็นโปรแกรมที่นักวิจัยส่วนใหญ่มักใช้จัดเก็บข้อมูลหรือวิเคราะห์ข้อมูลทางสถิติเบื้องต้นนั้น โปรแกรม SPSS จะให้เลือกว่าจะเปิด Sheet ใดที่ต้องการเปิด (เนื่องจากในโปรแกรม Excel ผู้ใช้งานอาจมีข้อมูลหลาย Sheet ใน 1 ไฟล์ข้อมูล)

ตัวอย่างการเปิดไฟล์ข้อมูลจากโปรแกรม Excel ในโปรแกรม SPSS

ยกตัวอย่างเช่น กรณีต้องการเปิดไฟล์ข้อมูลจากโปรแกรม Excel ชื่อ **Drying Temp.xlsx** ดังภาพที่ 3.19 ซึ่งไฟล์อยู่ใน **DATA (D:) > STAT DATA** และข้อมูลที่ต้องการป้อนในโปรแกรม SPSS อยู่ใน sheet 1 ดังภาพที่ 3.20 เมื่อพิจารณาจะพบว่าข้อมูลประกอบไปด้วย 4 คอลัมน์ (A, B, C และ D) และ 10 แถว (แถวที่ 1 คือ ชื่อตัวแปร ได้แก่ **TRT, REP, Drying Temp, L*** และแถวที่ 2-10 คือ ข้อมูลที่ต้องการเปิดในโปรแกรม SPSS)



ภาพที่ 3.19 ไฟล์ชื่อ Drying Temp.xlsx อยู่ใน DATA (D:) > STAT DATA

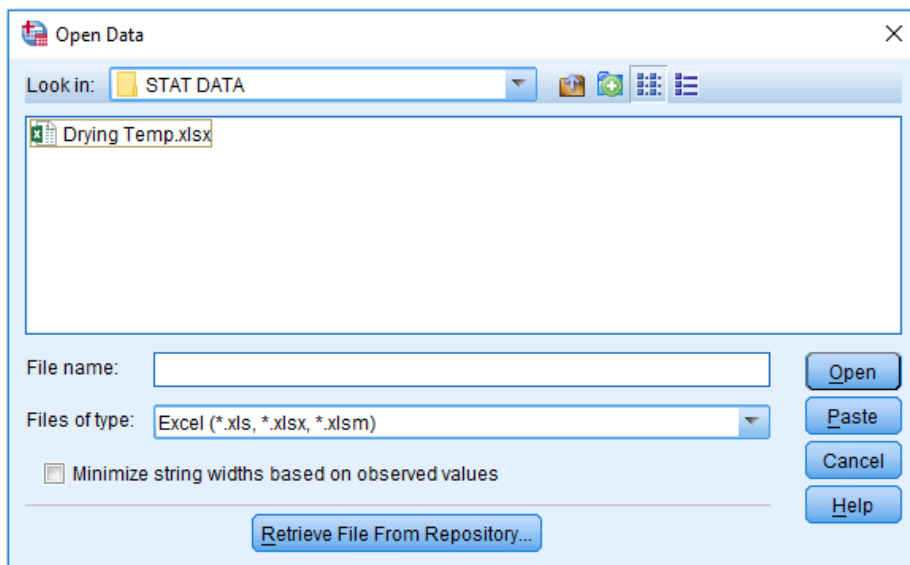
A screenshot of the Microsoft Excel application window titled 'Drying Temp.xlsx - Excel'. The ribbon shows 'Home', 'Insert', 'Page Layout', 'Formulas', 'Data', 'Review', 'View', and 'Help'. The formula bar shows 'E13'. The worksheet 'Sheet1' contains the following data:

	A	B	C	D	E	F	G
1	TRT	REP	Drying Temp	L*			
2	1	1	50	85.56			
3	1	2	50	85.44			
4	1	3	50	85.23			
5	2	1	60	80.36			
6	2	2	60	80.79			
7	2	3	60	80.11			
8	3	1	70	75.49			
9	3	2	70	75.88			
10	3	3	70	75.03			
11							

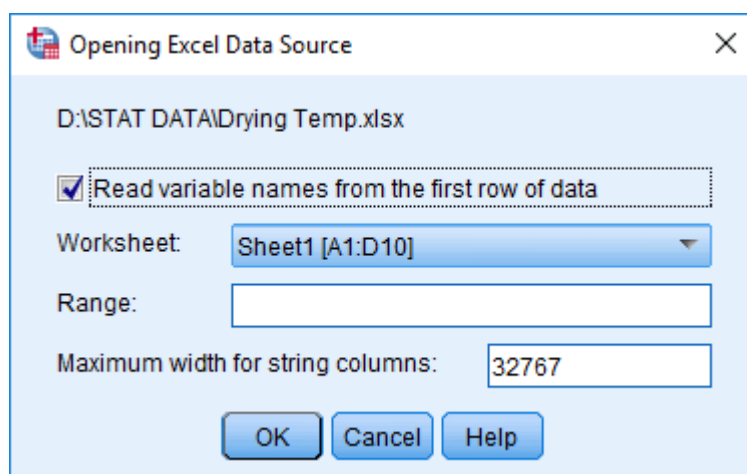
ภาพที่ 3.20 ตัวอย่างข้อมูลในโปรแกรม Excel ที่ต้องการเปิดในโปรแกรม SPSS

ขั้นตอนการเปิดไฟล์ Excel ในโปรแกรม SPSS ให้เปิดไฟล์ได้จากเมนู **File > Open > Data** จะปรากฏกล่อง Open Data ดังภาพที่ 3.21

- ที่ช่อง **Files of type** ให้เลือกชนิดของไฟล์ที่ต้องการเปิด ในที่นี้คือ Excel (*.xls, *.xlsx, *.xlsm)
- ที่ช่อง **Look in** ให้เลือก Drive และ Folder ที่มีไฟล์ชื่อ Drying Temp.xlsx อยู่ในกรณีนี้คือ DATA (D:) และ Folder ชื่อ STAT DATA
- กดเลือกไฟล์ชื่อ Drying Temp.xlsx
- แล้วกดคลิก Open จะปรากฏกล่อง Opening Excel Data Source ดังภาพที่ 3.22



ภาพที่ 3.21 การเลือกเปิดไฟล์ Excel ในโปรแกรม SPSS



ภาพที่ 3.22 การกำหนด Sheet และช่วงของข้อมูลในไฟล์ Excel ที่ต้องการเปิดใช้งานในโปรแกรม SPSS

- ในภาพที่ 3.22 จะปรากฏกล่อง Opening Excel Data Source ซึ่งโปรแกรม SPSS จะให้ระบุ Sheet และช่วงของข้อมูลที่ต้องการจะเปิด เมื่อพิจารณาจากภาพที่ 3.20 แสดงให้เห็นว่าช่วงของข้อมูลใน Excel อยู่ระหว่างตำแหน่งที่ A1 และ ตำแหน่งที่ D10 (ช่วงข้อมูลที่รวมชื่อตัวแปรและข้อมูลที่ต้องการวิเคราะห์) ในที่นี้แม้โปรแกรม SPSS จะระบุให้โดยอัตโนมัติแต่ผู้ใช้งานควรตรวจสอบความถูกต้องทุกครั้ง

- คำสั่ง ☒ **Read variable names from the first row of data** คือ การกำหนดให้โปรแกรม SPSS ตั้งชื่อของตัวแปรต่างๆ จากแถวแรกของข้อมูลในตาราง Excel ในกรณีที่ไม่ได้ตั้งชื่อตัวแปรไว้ในแถวแรก ไม่ต้องกดเลือก ☐ **Read variable names from the first row of data**

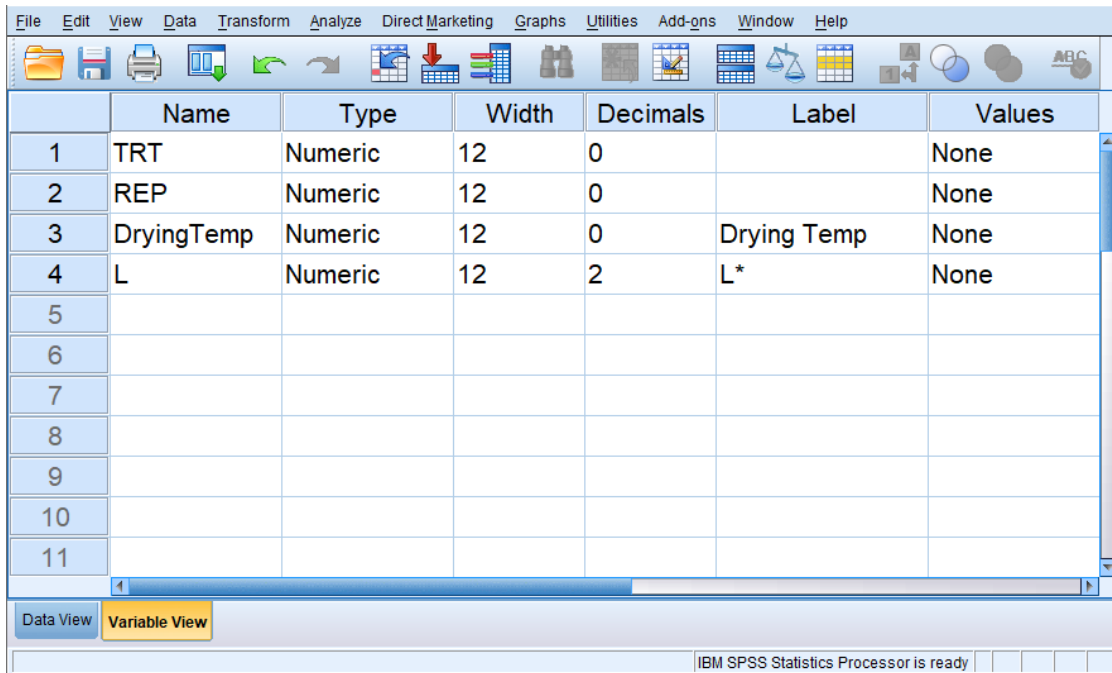
- เมื่อทำการเลือกเปิดไฟล์จากโปรแกรม Excel เสร็จเรียบร้อยแล้ว จะปรากฏหน้าต่าง Data View ดังภาพที่ 3.23 และหน้าต่าง Variable View ดังภาพที่ 3.24 เมื่อพิจารณาในหน้าต่าง Data View จะพบข้อมูลเช่นเดียวกับในไฟล์ Excel โดยไม่ต้องพิมพ์ชื่อตัวแปรเอง ทั้งนี้ให้สังเกตชื่อตัวแปร **DryingTemp** และ ตัวแปร **L** ในหน้าต่าง Data View ดังนี้

(1) ตัวแปร DryingTemp ถูกเปลี่ยนมาจากชื่อ Drying Temp ในโปรแกรม Excel เนื่องจากการเว้นวรรคซึ่งผิดกฎการตั้งชื่อตัวแปรในโปรแกรม SPSS ทั้งนี้เมื่อพิจารณาในหน้าต่าง Variable View จะพบว่าตัวแปร DryingTemp ได้ถูกกำหนดความหมายที่แท้จริงใน Cell คอลัมน์ Label เป็น Drying Temp โดยอัตโนมัติ

(2) ตัวแปร L ถูกเปลี่ยนมาจากชื่อ L* ในโปรแกรม Excel เนื่องจากการใช้เครื่องหมาย * ซึ่งผิดกฎการตั้งชื่อตัวแปรในโปรแกรม SPSS และเมื่อพิจารณาในหน้าต่าง Variable View จะพบว่าตัวแปร L ได้ถูกกำหนดความหมายที่แท้จริงใน Cell คอลัมน์ Label เป็น L* โดยอัตโนมัติ

	TRT	REP	DryingTemp	L	var	var	var
1	1	1	50	85.56			
2	1	2	50	85.44			
3	1	3	50	85.23			
4	2	1	60	80.36			
5	2	2	60	80.79			
6	2	3	60	80.11			
7	3	1	70	75.49			
8	3	2	70	75.88			
9	3	3	70	75.03			
10							

ภาพที่ 3.23 ตัวอย่างหน้าต่าง Data View ที่ได้จากการเปิดข้อมูลจากไฟล์ Excel



ภาพที่ 3.24 ตัวอย่างหน้าต่าง Variable View ที่ได้จากการเปิดข้อมูลจากไฟล์ Excel

3.3.2 การป้อนข้อมูลโดยตรง

การป้อนข้อมูลโดยตรงในโปรแกรม SPSS สามารถทำได้โดยตรงในโปรแกรม SPSS หรือสามารถไปป้อนข้อมูลในโปรแกรม Excel หรือโปรแกรมที่มีลักษณะของข้อมูลเป็นแบบ Spread sheet เพื่อความสะดวกและง่ายในการป้อนข้อมูล จากนั้น Copy ข้อมูลในโปรแกรกดังกล่าวแล้วนำมา Paste วางลงในโปรแกรม SPSS อย่างไรก็ตามแม้ว่าจะสะดวกในการป้อนข้อมูลแต่ผู้ใช้งานยังต้องมาพิมพ์ข้อมูลรายละเอียดของตัวแปรเองในหน้าต่าง Variable View

ในการป้อนข้อมูลโดยตรงในโปรแกรม SPSS นักวิจัยต้องพิจารณาข้อมูลที่มีอยู่เป็นครั้งแรกแล้วตอบคำถามเพื่อกำหนดรายละเอียดข้อมูลของตัวแปรในหน้าต่าง Variable View เช่น

- (1) จำนวนตัวแปร var ที่ต้องใช้ในการป้อนข้อมูล เท่ากับ.....
- (2) ชื่อของตัวแปรแต่ละตัว (Name) คือ.....
- (3) จำนวนทศนิยมของข้อมูลแต่ละตัวแปร (Decimals) คือ.....
- (4) ระบุความหมายที่แท้จริงของตัวแปร (Label) ☐ ไม่ระบุ ☐ ระบุ คือ.....
- (5) กำหนดค่าให้ตัวแปร (Values) ☐ ไม่กำหนด ☐ กำหนด คือ.....

ตัวอย่างการป้อนข้อมูลเองโดยตรงในโปรแกรม SPSS

นักพัฒนาผลิตภัณฑ์ต้องการแปรรูปขนุนแผ่น จึงนำเนื้อขนุนมาป่นกับส่วนผสมแล้วนำไปอบแห้งด้วยเครื่อง Tray dryer โดยใช้อุณหภูมิที่ต่างกัน 3 ระดับ ได้แก่ 50, 60 และ 70 องศาเซลเซียส ใช้เวลาในการอบเท่ากัน นำขนุนแผ่นไปวัดค่าความสว่าง (L^*) ด้วยเครื่องวัดค่าสี Chroma meter ตัวอย่างละ 3 ซ้ำ ได้ผลแสดงดังตารางที่ 3.1

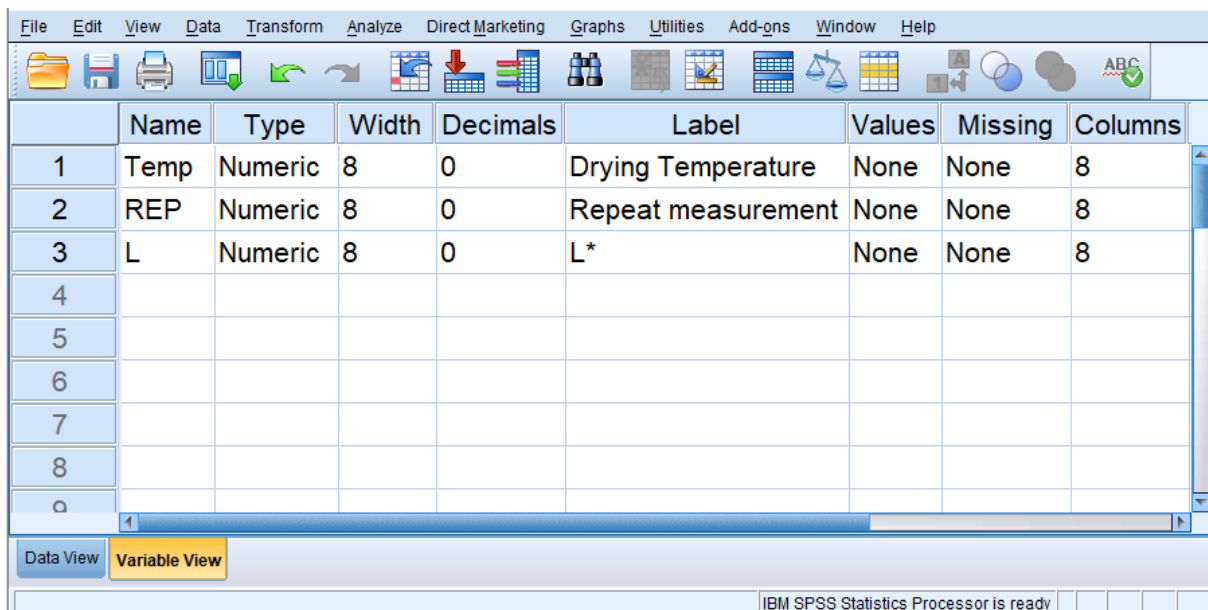
ตารางที่ 3.1 ค่าความสว่าง (L^*) ของขนุนแผ่นเมื่ออบด้วยอุณหภูมิที่ต่างกัน

อุณหภูมิในการอบ ($^{\circ}\text{C}$)	การวัดค่าซ้ำที่	ค่าความสว่าง (L^*)
50	1	98
	2	96
	3	95
60	1	85
	2	88
	3	86
70	1	75
	2	79
	3	74

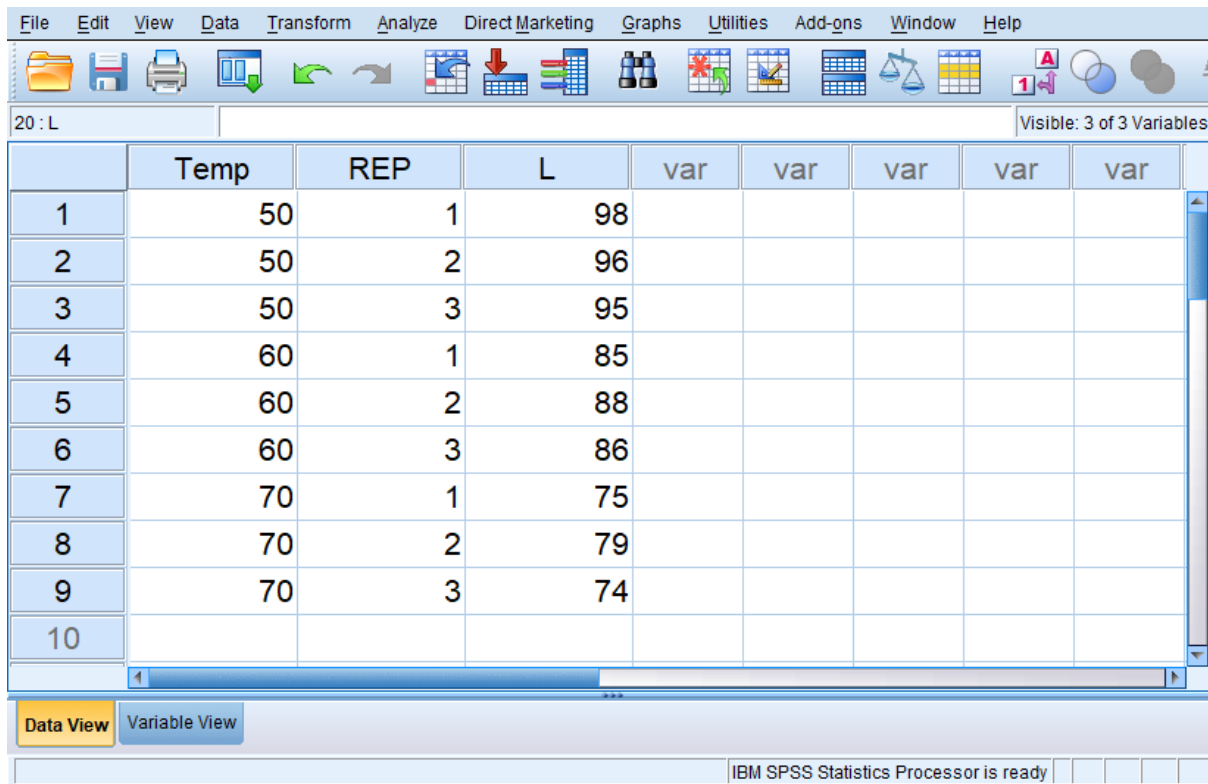
จากตารางที่ 3.1 พิจารณาข้อมูลเพื่อกำหนดรายละเอียดของตัวแปรในหน้าต่าง Variable View ดังนี้

- (1) จำนวนตัวแปร var ที่ต้องใช้ในการป้อนข้อมูล เท่ากับ 3 ตัว (อุณหภูมิในการอบ, การวัดค่าซ้ำ และ ค่าความสว่าง L^*)
- (2) ชื่อของตัวแปรแต่ละตัว (Name) คือ Temp, Rep, L
- (3) จำนวนทศนิยมของข้อมูลแต่ละตัวแปร (Decimals) คือ 0 (เนื่องจากตัวแปร Temp และ Rep เป็นข้อมูลชนิด Nominal และเมื่อพิจารณาตัวแปร L แม้ว่าจะเป็นข้อมูลชนิด Scale แต่จากข้อมูลดิบพบว่าเป็นเลขจำนวนเต็ม จึงกำหนดให้จำนวนทศนิยมเป็น 0 เช่นเดียวกัน)
- (4) ระบุความหมายที่แท้จริงของตัวแปร (Label) ☐ ไม่ระบุ ☒ ระบุ คือ
Drying Temperature.....สำหรับตัวแปร Temp.....
Repeat measurement.....สำหรับตัวแปร Rep.....
 L^*สำหรับตัวแปร L.....
- (5) กำหนดค่าให้ตัวแปร (Values) ☒ ไม่กำหนด ☐ กำหนด คือ.....

เมื่อพิจารณาข้อมูลแล้ว ขั้นตอนต่อไป คือ กำหนดรายละเอียดให้ตัวแปรแต่ละตัวในหน้าต่าง Variable View (ภาพที่ 3.25) และป้อนข้อมูลในหน้าต่าง Data View (ภาพที่ 3.26)



ภาพที่ 3.25 การกำหนดรายละเอียดให้ตัวแปรเมื่อต้องป้อนข้อมูลเองจากตารางที่ 3.1



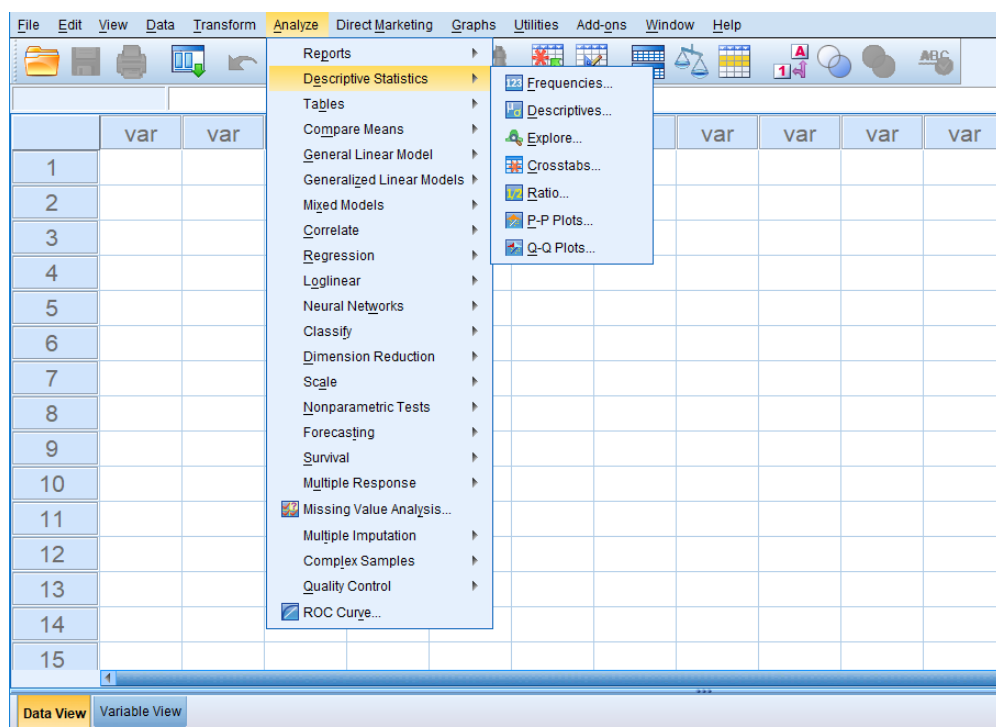
ภาพที่ 3.26 การป้อนข้อมูลเองในโปรแกรม SPSS โดยใช้ข้อมูลดิบจากตารางที่ 3.1

3.4 เมนูสำหรับการวิเคราะห์ผลทางสถิติ

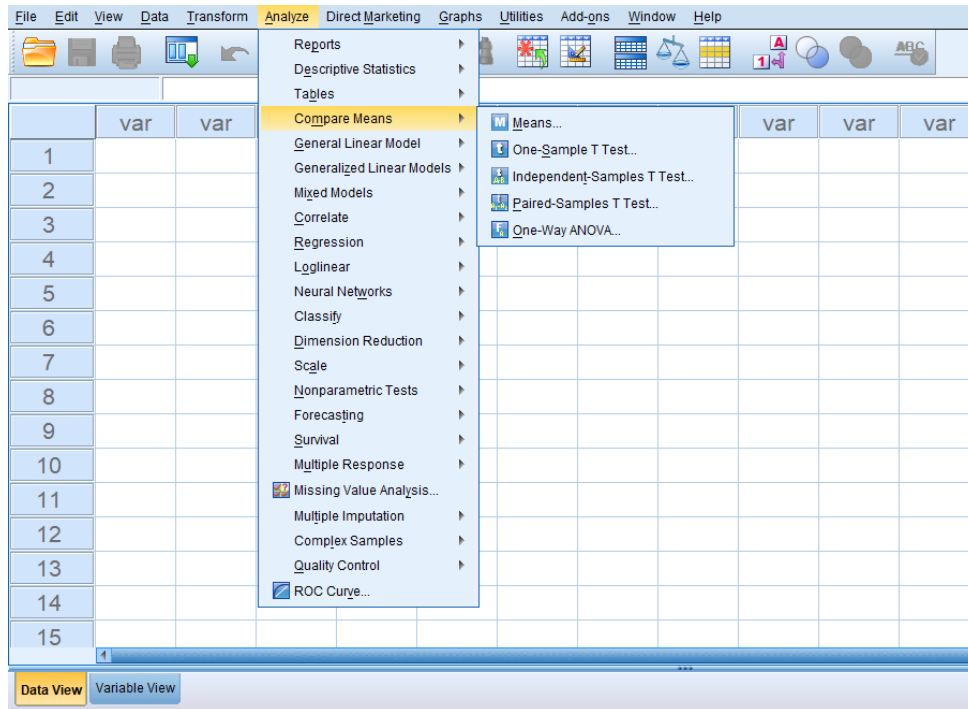
ในการวิเคราะห์ผลทางสถิติจะใช้เมนูคำสั่ง **Analyze** ด้านบนจอ เมื่อกดคลิกที่ Analyze จะปรากฏเมนูย่อยที่ใช้ในการวิเคราะห์ทางสถิติต่างๆ อาทิเช่น

- **Descriptive Statistics** (ภาพที่ 3.27) ใช้ในการคำนวณค่าเฉลี่ย ค่าส่วนเบี่ยงเบนมาตรฐาน ค่าต่ำสุดและค่าสูงสุดของข้อมูลที่เป็นตัวเลข
- **Compare Means** (ภาพที่ 3.28) ใช้ในการเปรียบเทียบความแตกต่างของค่าเฉลี่ย
- **General Linear Model** (ภาพที่ 3.29) ใช้ในการวิเคราะห์ความแปรปรวนของข้อมูล
- **Regression** (ภาพที่ 3.30) ใช้ในการหาความสัมพันธ์ระหว่างตัวแปรต้นกับตัวแปรตามโดยการสร้างสมการถดถอย

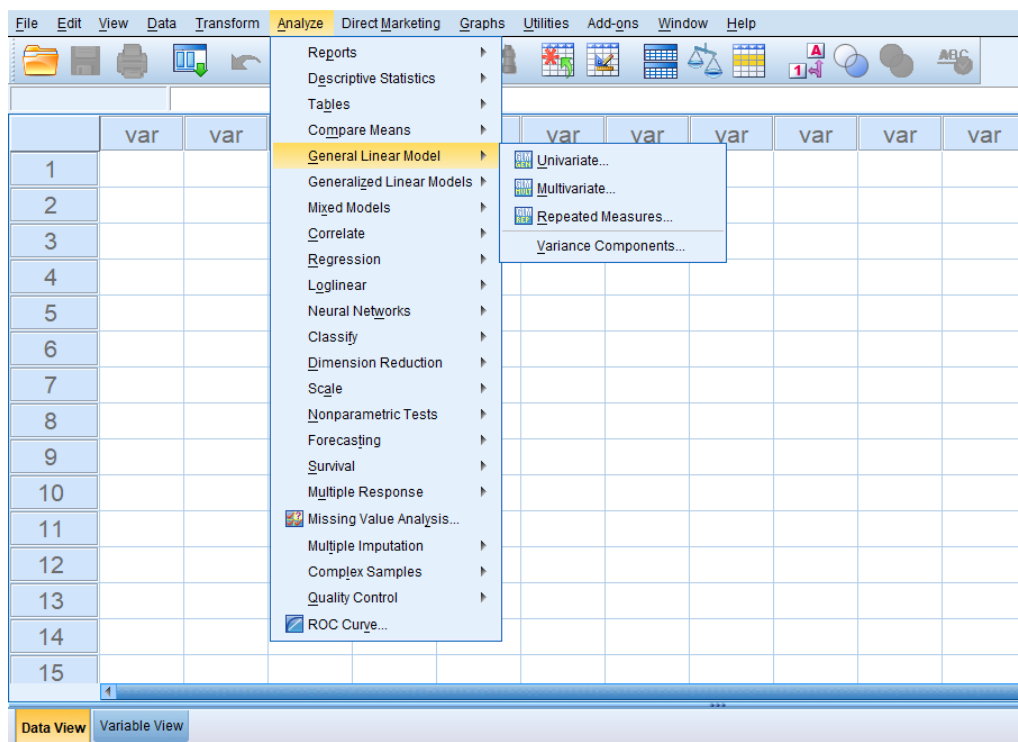
รายละเอียดการใช้แต่ละเมนูคำสั่งรวมทั้งการป้อนข้อมูลของแต่ละวิธีวิเคราะห์ และการแปลความหมายที่ได้จากการวิเคราะห์ค่าทางสถิติจะกล่าวถึงในบทถัดไป



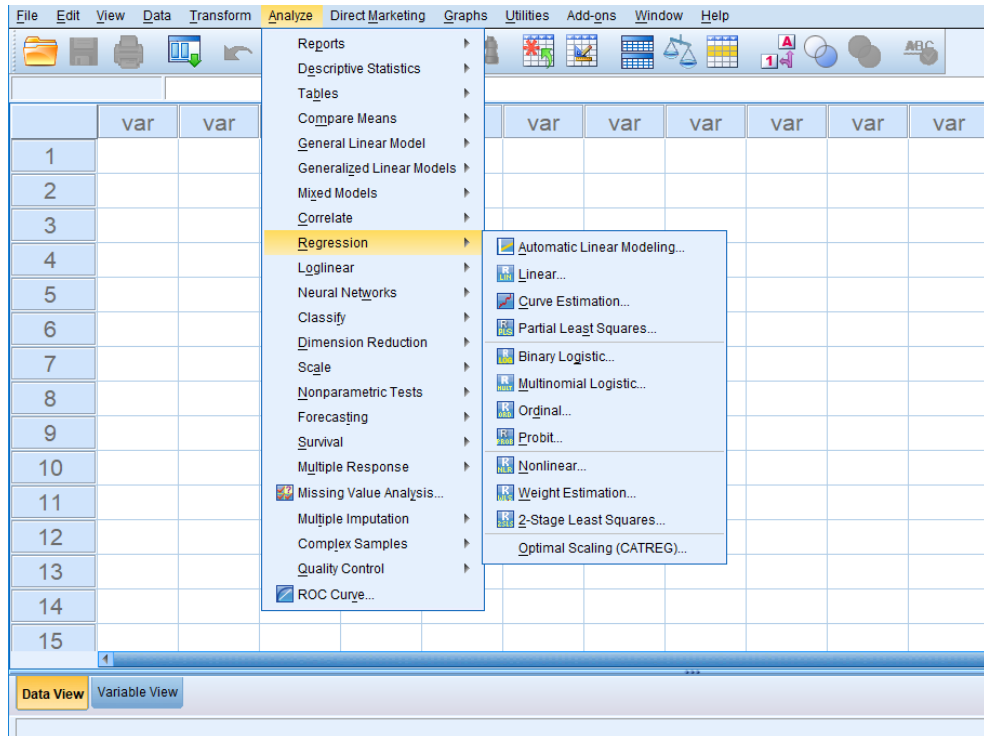
ภาพที่ 3.27 เมนูคำสั่ง Analyze ➤ Descriptive Statistics



ภาพที่ 3.28 เมนูคำสั่ง Analyze ➤ Compare Means



ภาพที่ 3.29 เมนูคำสั่ง Analyze ➤ General Linear Model

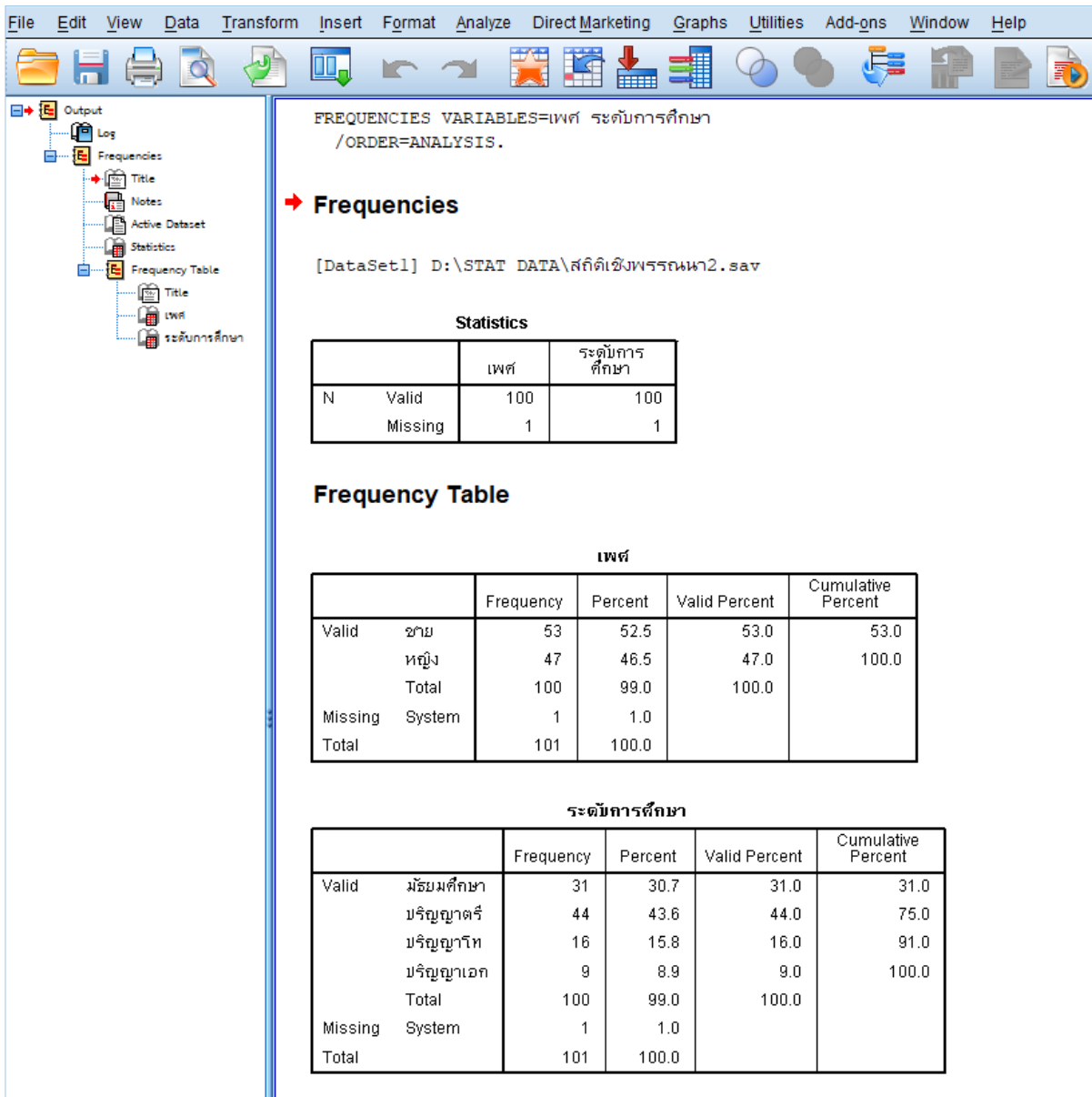


ภาพที่ 3.30 เมนูคำสั่ง Analyze ➤ Regression

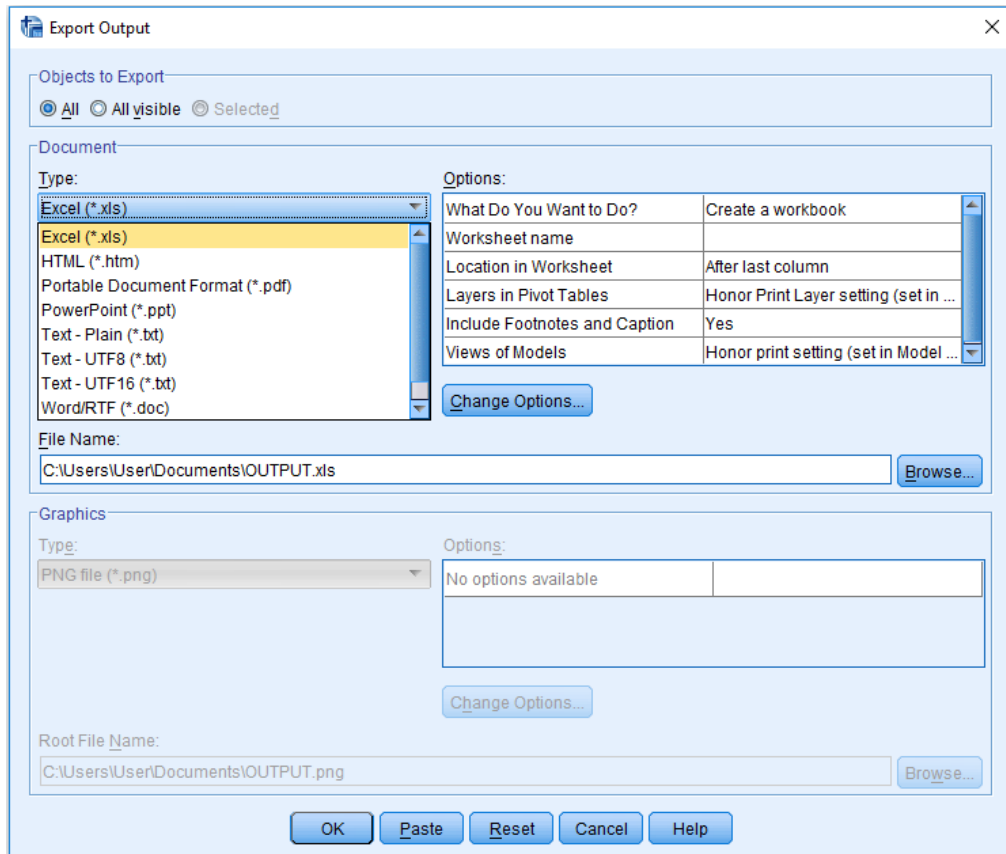
3.5 การบันทึกผลลัพธ์ (Output) ของโปรแกรม SPSS

ผลลัพธ์จากการวิเคราะห์ข้อมูลทางสถิติด้วยโปรแกรม SPSS จะแสดงในรูปแบบตาราง Pivot ประกอบด้วยส่วนต่างๆ อาทิเช่น ชื่อวิธีวิเคราะห์ทางสถิติ ชื่อไฟล์ของข้อมูลที่จะวิเคราะห์ และข้อมูลผลการวิเคราะห์ในตาราง ดังภาพที่ 3.31 ซึ่งนักวิจัยสามารถบันทึกผลลัพธ์ที่ได้ด้วยคำสั่ง **Save** หรือ **Save As** โดยผลลัพธ์นั้นจะถูกบันทึกไว้ในรูปแบบ *.spv ซึ่งไม่สามารถนำไปเปิดด้วยโปรแกรมอื่นได้ นอกจากนี้นักวิจัยยังสามารถพิมพ์รายละเอียดเพิ่มเติมได้ในส่วนของชื่อบนตารางโดยการกดดับเบิลคลิกแล้วพิมพ์ข้อความที่ต้องการเพิ่มเติม

นักวิจัยสามารถบันทึกข้อมูลเพื่อนำไปใช้กับโปรแกรมอื่นได้ เช่น Word หรือ Excel โดยใช้เมนูคำสั่ง **File ➤ Export** และเลือกลักษณะของผลลัพธ์ที่ต้องการที่เมนู **Type:** ดังภาพที่ 3.32 แล้วพิมพ์ชื่อไฟล์ที่ต้องการในช่อง **File Name:** จากนั้นกดคลิก OK จะได้ผลลัพธ์ในรูปแบบไฟล์ที่ต้องการ หรือ สามารถคลิกเมาส์ปุ่มขวาในตาราง Output ที่สนใจ จากนั้นกด Copy แล้วไปกด Paste วางในโปรแกรมไฟล์ Word, PowerPoint และ Excel ได้ โดยนักวิจัยสามารถเลือกชนิดของการ Paste วางได้หลายแบบ เช่น รูปภาพ หรือ ข้อมูลในรูปแบบตาราง



ภาพที่ 3.31 ตัวอย่างผลลัพธ์ที่ได้จากการวิเคราะห์ด้วยคำสั่ง Frequencies



ภาพที่ 3.32 เมนูคำสั่ง Export Output



บทที่ 4

สถิติเชิงพรรณนาที่ใช้ในงานวิจัย

สถิติเชิงพรรณนา (Descriptive statistics) เป็นสถิติที่ใช้สรุปลักษณะของกลุ่มข้อมูล โดยเป็นข้อมูลที่เกิดขึ้นจากตัวอย่างหรือประชากร สถิติเชิงพรรณนา เช่น จำนวน ร้อยละ ค่าเฉลี่ย และค่าส่วนเบี่ยงเบนมาตรฐาน ในขั้นเริ่มต้นของงานพัฒนาผลิตภัณฑ์มักใช้แบบสอบถามในการสำรวจความคิดเห็น พฤติกรรม และความต้องการของผู้บริโภคที่มีต่อผลิตภัณฑ์ที่ต้องการพัฒนา แล้วนำข้อมูลมาวิเคราะห์เพื่อใช้ในการวิเคราะห์โอกาสในการพัฒนาผลิตภัณฑ์และสร้างแนวความคิดผลิตภัณฑ์ใหม่ วิเคราะห์จุดบกพร่องของผลิตภัณฑ์ที่ผู้บริโภคไม่พึงประสงค์ พิจารณาปัจจัยที่มีผลต่อการยอมรับและตัดสินใจซื้อของผู้บริโภค หรือ การวิเคราะห์หา กลุ่มเป้าหมายที่คาดว่าจะยอมรับผลิตภัณฑ์ที่จะพัฒนาขึ้นเพื่อใช้ในการกำหนดกลุ่มเป้าหมายและส่งเสริมทางด้านการตลาด เมื่อสิ้นสุดการพัฒนาผลิตภัณฑ์ได้สูตรที่เหมาะสมแล้วนักวิจัยต้องทำการทดสอบการยอมรับของผู้บริโภคเพื่อประเมินโอกาสที่ผลิตภัณฑ์จะได้รับการยอมรับเมื่อทำการวางจำหน่าย

การสำรวจพฤติกรรมและการทดสอบการยอมรับของผู้บริโภคต้องใช้ผู้ทดสอบจำนวนมากโดยใช้แบบสอบถามเป็นเครื่องมือในการวิจัย ดังนั้นนักวิจัยต้องมีความรู้ในการสร้างแบบสอบถาม และสามารถเลือกใช้สถิติเชิงพรรณนาเพื่อสรุปลักษณะของข้อมูลได้อย่างเหมาะสม เช่น ผู้ตอบแบบสอบถามกลุ่มใดที่สนใจในตัวผลิตภัณฑ์ใหม่ กลุ่มผู้บริโภคให้ความสำคัญต่อคุณลักษณะใดของผลิตภัณฑ์ เพื่อนำข้อมูลเหล่านี้ไปใช้ประกอบการวางแผนการวิจัยและพัฒนาผลิตภัณฑ์ต่อไป

4.1 ประเภทของสถิติเชิงพรรณนา

การเลือกใช้สถิติเชิงพรรณนาวิธีใดนั้นขึ้นอยู่กับประเภทของข้อมูลหรือสเกลที่ใช้ในการรวบรวมข้อมูล ในงานวิจัยเรื่องหนึ่งๆ อาจจะประกอบด้วยข้อมูลหลายชนิดที่มีทั้งข้อมูลเชิงคุณภาพ (สเกลนามกำหนด และ สเกลอันดับ) และข้อมูลเชิงปริมาณ (สเกลอัตราภาค และ สเกลอัตราส่วน) ดังนั้นการเลือกใช้สถิติเชิงพรรณนาจึงแบ่งออกได้เป็น 2 กลุ่ม คือ

4.1.1 สถิติเชิงพรรณนาสำหรับข้อมูลเชิงคุณภาพหรือเชิงกลุ่ม

ข้อมูลเชิงคุณภาพ (Qualitative data) บางครั้งเรียกว่า ข้อมูลเชิงกลุ่ม หรือข้อมูลจำแนกกลุ่ม เป็นข้อมูลที่อยู่ในรูปข้อความและแสดงลักษณะที่แตกต่างกัน ข้อมูลที่ได้จากสเกลนามกำหนด (Nominal scale)

และสเกลอันดับ (Ordinal scale) จัดเป็นข้อมูลเชิงคุณภาพ สถิติเชิงพรรณนาที่ใช้วิเคราะห์ข้อมูลเชิงคุณภาพ ได้แก่ ความถี่หรือจำนวน ร้อยละ และ ฐานนิยม

4.1.2 สถิติเชิงพรรณนาสำหรับข้อมูลเชิงปริมาณ

ข้อมูลเชิงปริมาณ (Quantitative data) เป็นข้อมูลตัวเลขที่สามารถวัดค่าได้ว่ามีค่ามากหรือน้อย เช่น รายได้ อายุ เกรดเฉลี่ยสะสม ยอดขาย ปริมาณไขมัน ปริมาณโปรตีน ราคาขาย ระดับความพึงพอใจ และ คะแนนความชอบ ข้อมูลที่ได้จากสเกลอันดับ (Interval scale) และสเกลอัตราส่วน (Ratio scale) จัดเป็นข้อมูลเชิงปริมาณซึ่งจะต้องคำนวณค่าสถิติเพื่อสรุปลักษณะข้อมูล สถิติเชิงพรรณนาที่ใช้วิเคราะห์ข้อมูลเชิงปริมาณ ได้แก่ ค่ากลาง (ค่าเฉลี่ย, ค่ามัธยฐาน และ ค่าฐานนิยม) และค่าการกระจายของข้อมูล (ค่าพิสัย, ค่าแปรปรวน และ ค่าส่วนเบี่ยงเบนมาตรฐาน)

4.2 การเก็บรวบรวมข้อมูลเพื่อใช้ในการวิเคราะห์สถิติเชิงพรรณนา

ในงานพัฒนาผลิตภัณฑ์เมื่อนักวิจัยเก็บรวบรวมข้อมูลของผู้บริโภคโดยใช้แบบสอบถาม ข้อมูลที่ได้ถือได้ว่าเป็นข้อมูลปฐมภูมิ โดยแบบสอบถาม คือ เครื่องมือที่ใช้ในการเก็บรวบรวมข้อมูล ซึ่งอาจทำได้ทั้งการสัมภาษณ์แบบส่วนตัว สัมภาษณ์ทางโทรศัพท์ สัมภาษณ์ด้วยระบบออนไลน์ และส่งทางไปรษณีย์ คำถามที่มักถูกเลือกใช้ในแบบสอบถาม เช่น

- (1) คำถามด้านประชากรศาสตร์/สังคมเศรษฐกิจ เช่น อายุ เพศ ระดับการศึกษา อาชีพ สถานภาพ การสมรส รายได้ โดยนักวิจัยใช้ข้อมูลเหล่านี้เพื่อพิจารณาลักษณะของผู้ตอบแบบสอบถาม
- (2) คำถามด้านทัศนคติ/ความคิดเห็น เช่น ความพึงพอใจ แนวโน้ม หรือความรู้สึกต่อสิ่งใดสิ่งหนึ่ง การแสดงความคิดเห็นในเรื่องใดเรื่องหนึ่ง ทั้งทัศนคติและความคิดเห็นสามารถใช้แทนกันได้
- (3) คำถามเกี่ยวกับพฤติกรรมและแรงจูงใจ เช่น แรงจูงใจในการซื้อผลิตภัณฑ์ พฤติกรรมการซื้อ และการใช้ผลิตภัณฑ์ ซื้ออะไร ซื้อเท่าไร ซื้ออย่างไร ซื้อที่ไหน ซื้อเมื่อใด ใครเป็นผู้ซื้อ

รูปแบบของคำถามในแบบสอบถาม แบ่งออกได้เป็น 2 ประเภท ดังนี้

4.2.1 คำถามปลายเปิด (Open-ended questions) หมายถึง คำถามซึ่งเปิดโอกาสให้ผู้ตอบแบบสอบถามตอบได้อย่างอิสระโดยไม่มีขีดจำกัด แบ่งออกได้เป็น 7 แบบ ได้แก่

- (1) **คำถามแบบไม่มีโครงสร้าง (Unstructured questions)** เป็นคำถามที่ไม่มีการวางแผนหรือจัดลำดับการตอบเอาไว้ ตัวอย่างเช่น

ท่านมีความคิดเห็นอย่างไรเกี่ยวกับไอศกรีมข้าวเหนียวมะม่วงที่ท่านชิม.....

(2) คำถามแบบมีโครงสร้าง (Structured questions) เป็นคำถามที่มีการวางแผนหรือจัดลำดับการตอบเอาไว้อย่างแน่นอนเพื่อให้ผู้ตอบแบบสอบถามตอบตามลำดับในแต่ละข้อ ตัวอย่างเช่น

ท่านชอบรับประทานข้าวเหนียวมะม่วงหรือไม่ เพราะเหตุใด

☐ ชอบ เพราะ

☐ ไม่ชอบ เพราะ

ท่านเคยปอกมะม่วงสุกรับประทานเองหรือไม่

ปัญหาที่ท่านพบในการปอกมะม่วงสุกรับประทานเองคืออะไร.....

(3) คำถามแบบโยงความสัมพันธ์ระหว่างคำพูดหรือข้อความ (Word association) วิธีนี้มีประโยชน์ในการให้ผู้ตอบแบบสอบถามกล่าวถึงชื่อตราสินค้า โลโก้ และคำขวัญ หรือคำพูดที่อยู่ในใจเมื่อได้อ่านข้อความนั้น ตัวอย่างเช่น

ยี่ห้อขนมเบเกอรี่ที่คนนิยมมากที่สุด คือ

โลโก้ของร้านไก่ทอดเคเอฟซี คือ

(4) การเติมประโยคให้สมบูรณ์ (Sentence completion) เป็นเทคนิคการสร้างภาพซึ่งผู้ตอบแบบสอบถามจะต้องเติมประโยคให้สมบูรณ์ที่เกิดขึ้นจากจิตใจ ตัวอย่างเช่น

รูปแบบบรรจุภัณฑ์ขนมกล้วยที่ท่านชอบมากที่สุด คือ

ไส้ชาลาเปาที่ท่านชอบรับประทานมากที่สุด คือ.....

(5) การเติมเรื่องราวให้สมบูรณ์ (Story completion) เป็นเรื่องราวที่ยังไม่สมบูรณ์ ต้องการให้ผู้ตอบแบบสอบถามเติมข้อความเพื่อให้เป็นเรื่องราวที่สมบูรณ์ ตัวอย่างเช่น

เมื่อท่านรับประทานขนมกล้วยไม่หมด ท่านจะ.....

เมื่อท่านไปซื้อน้ำอัดลมยี่ห้อประจำแล้วพบว่าสินค้าหมด ท่านจะ.....

(6) การเติมภาพให้สมบูรณ์ (Picture completion) เป็นภาพที่นำเสนอไว้ ซึ่งผู้ตอบแบบสอบถามจะต้องเติมภาพหรือข้อความที่ว่างเอาไว้ ตัวอย่างเช่น



โลโก้ของร้านสะดวกซื้อ 7-11 ที่ขาดหายไป คือ

(7) Thematic Apperception Test (TAT) เป็นภาพที่นำเสนอไว้ให้ผู้ตอบแบบสอบถามสร้างเรื่องราวที่เกี่ยวข้องกับภาพนั้นตามความคิดเห็นของตนเอง ตัวอย่างเช่น



โลโก้ของร้านกาแฟสตาร์บัคมีลักษณะ.....

4.2.2 คำถามปลายปิด (Close-ended questions) หมายถึง คำถามซึ่งประกอบไปด้วยคำตอบที่มีโครงสร้าง ผู้ตอบแบบสอบถามต้องเลือกคำตอบจากกลุ่มของคำตอบที่มีให้เลือก มีได้หลายแบบ ดังนี้

(1) คำถามแบบมีคำตอบให้เลือก 2 ข้อ (Two-way questions) เป็นคำถามซึ่งมีคำตอบ 2 คำตอบ ตัวอย่างเช่น

ท่านเคยรับประทานมะม่วงน้ำดอกไม้หรือไม่ ☐ เคย ☐ ไม่เคย

(2) คำถามแบบหลายตัวเลือก (Multiple choice questions) เป็นคำถามที่มีหลายตัวเลือกโดยให้ผู้ตอบแบบสอบถามเลือกหนึ่งคำตอบจากตัวเลือกทั้งหมด ตัวอย่างเช่น

ท่านซื้อมะม่วงน้ำดอกไม้สุกจากสถานที่ใดบ่อยที่สุด

- | | | |
|---------------------------------|-----------------------------------|--|
| ☐ ตลาดสด | ☐ ตลาดนัด | ☐ ซูเปอร์มาร์เก็ต |
| ☐ รถเข็น | ☐ โรงอาหาร | ☐ ร้านขายของฝาก |

(3) คำถามแบบให้ตอบได้หลายข้อ (Checklists) เป็นคำถามที่มีหลายตัวเลือกโดยให้ผู้ตอบแบบสอบถามเลือกคำตอบได้มากกว่า 1 ข้อ ตัวอย่างเช่น

ท่านเคยรับประทานผลิตภัณฑ์จากมะม่วงสุกใดบ้าง (เลือกตอบได้มากกว่า 1 ข้อ)

- | | | | |
|---|---|---------------------------------------|-------------------------------------|
| ☐ ข้าวเหนียวมะม่วง | ☐ เค้กมะม่วง | ☐ มะม่วงแผ่น | ☐ วุ้นมะม่วง |
| ☐ ไอศกรีมมะม่วง | ☐ น้ำแข็งไสมะม่วง (빙수) | ☐ เยลลี่มะม่วง | ☐ น้ำมะม่วง |

(4) คำถามแบบจัดอันดับ (Ranking questions) เป็นคำถามที่ต้องการให้ผู้ตอบแบบสอบถามจัดอันดับคุณลักษณะที่สนใจโดยการเรียงลำดับ ตัวอย่างเช่น

ให้ท่านจัดอันดับความสำคัญ 3 อันดับแรกของปัจจัยที่มีผลต่อการเลือกซื้อข้าวเหนียวมะม่วงพร้อมรับประทาน (1 คือ สำคัญมากที่สุด)

- | | |
|---|---|
| ☐ ลักษณะปรากฏของมะม่วงสุก | ☐ ปริมาณข้าวเหนียว |
| ☐ ลักษณะปรากฏของข้าวเหนียว | ☐ ปริมาณมะม่วง |
| ☐ ราคา | ☐ ปริมาณน้ำกะทิ |
| ☐ บรรจุภัณฑ์ | ☐ สีของข้าวเหนียว |

(5) คำถามแบบใช้มาตรวัดของลิเคิร์ต (Likert scale) เป็นสเกลวัดทัศนคติของผู้ตอบแบบสอบถาม เช่น ความพึงพอใจหรือความไม่พอใจ การยอมรับหรือไม่ยอมรับ อาจกำหนดเป็นข้อความหรือตัวเลขแสดงระดับความคิดเห็นของผู้ตอบแบบสอบถาม ตัวอย่างเช่น

ท่านเห็นด้วยหรือไม่ว่าไอศกรีมมะม่วงน้ำดอกไม้วอร์มิกลีนของมะม่วงสุกจากธรรมชาติ

- ☐ เห็นด้วยอย่างยิ่ง ☐ เห็นด้วย ☐ ไม่แน่ใจ ☐ ไม่เห็นด้วย ☐ ไม่เห็นด้วยอย่างยิ่ง

(6) คำถามแบบใช้มาตรวัดเจตคติ (Semantic differential) เป็นสเกลการวัดลักษณะทัศนคติ 7 ระดับ โดยแบ่งเป็นช่วงๆ ตั้งแต่ด้านซ้ายสุดของสเกลหมายถึงลักษณะที่ดีมาก ด้านขวาสุดหมายถึงลักษณะที่แย่มาก มักใช้ในการวัดคุณสมบัติด้านจุดแข็ง/จุดอ่อน หรือคุณสมบัติดี/ไม่ดี ตัวอย่างเช่น

ท่านมีความคิดเห็นต่อไอศกรีมข้าวเหนียวมะม่วงสุกที่ท่านชิมต่อไปนี้อย่างไร

- | | | |
|-------------------------|------------------------------------|---------------------------|
| รสชาติอร่อย | ___: ___: ___: ___: ___: ___: ___: | รสชาติไม่อร่อย |
| กลิ่นมะม่วงชัดเจน | ___: ___: ___: ___: ___: ___: ___: | กลิ่นมะม่วงน้อย |
| ข้าวเหนียวนุ่ม | ___: ___: ___: ___: ___: ___: ___: | ข้าวเหนียวแข็ง |
| ละลายช้า | ___: ___: ___: ___: ___: ___: ___: | ละลายเร็ว |
| เนื้อสัมผัสเนียน | ___: ___: ___: ___: ___: ___: ___: | เนื้อสัมผัสไม่เนียน |
| เนื้อมะม่วงมีปริมาณพอดี | ___: ___: ___: ___: ___: ___: ___: | เนื้อมะม่วงมีปริมาณน้อยไป |

(7) คำถามแบบใช้สเกลความสำคัญ (Importance scale) เป็นคำถามที่ให้ผู้ตอบแบบสอบถามให้คะแนนความสำคัญของคุณสมบัติ ตัวอย่างเช่น

คุณสมบัติของไอศกรีมข้าวเหนียวมะม่วงต่อไปนี้มีความสำคัญสำหรับท่านมากน้อยเพียงใด	สำคัญอย่างยิ่ง	สำคัญมาก	สำคัญ	ไม่สำคัญ	ไม่สำคัญอย่างยิ่ง
รสชาติของมะม่วงจากธรรมชาติ	_____	_____	_____	_____	_____
สีของมะม่วงจากธรรมชาติ	_____	_____	_____	_____	_____
รสหวาน	_____	_____	_____	_____	_____
การละลาย	_____	_____	_____	_____	_____
เนื้อสัมผัสของไอศกรีม	_____	_____	_____	_____	_____
เนื้อสัมผัสของข้าวเหนียว	_____	_____	_____	_____	_____

(8) คำถามแบบใช้สเกลมาตราประมาณค่า (Rating scale) ใช้วัดคุณลักษณะต่างๆ ที่คนคิดที่ไม่สามารถวัดออกมาเป็นตัวเลขได้ แต่แสดงออกมาเป็นลักษณะต่างๆ ในระดับที่ต่างกัน เช่น ความชอบ ความพอใจ ตัวอย่างเช่น

ท่านชอบไอศกรีมข้าวเหนียวมะม่วงระดับใด				
ชอบมากที่สุด	ชอบมาก	ชอบปานกลาง	ไม่ชอบ	ไม่ชอบมากที่สุด
☐	☐	☐	☐	☐

ท่านพอใจต่อการบริการรถมอเตอร์ไซค์รับจ้างในม.เกษตรระดับใด				
พอใจมากที่สุด	พอใจมาก	พอใจปานกลาง	ไม่พอใจ	ไม่พอใจมากที่สุด
☐	☐	☐	☐	☐

(9) คำถามแบบใช้สเกลความตั้งใจซื้อ (Intention-to-buy scale) เป็นคำถามที่ให้ผู้ตอบแบบสอบถามแสดงความตั้งใจซื้อ ตัวอย่างเช่น

หากมีไอศกรีมข้าวเหนียวมะม่วงจำหน่ายในท้องตลาด ท่านจะซื้อหรือไม่		
☐ ซื้อแน่นอน	☐ ไม่แน่ใจ	☐ ไม่ซื้อแน่นอน

4.3 สถิติเชิงพรรณนาสำหรับข้อมูลเชิงคุณภาพหรือเชิงกลุ่ม

4.3.1 การกำหนดรหัสตัวแปรให้ข้อมูลเชิงคุณภาพ

ก่อนทำการวิเคราะห์ข้อมูลเชิงคุณภาพที่ได้จากสเกลนามกำหนดและสเกลอันดับ จะต้องทำการกำหนดรหัสคำตอบให้ค่าตัวแปรก่อน ในทางปฏิบัติเมื่อนักวิจัยสร้างแบบสอบถามจะทำการกำหนดรหัสตัวแปรควบคู่ไปด้วยเลย คำถาม 1 ข้อ จะสร้างเป็นตัวแปรได้อย่างน้อย 1 ตัว บางครั้งเราเห็นในแบบสอบถามด้านขวามือมีช่องสี่เหลี่ยมสำหรับให้เจ้าหน้าที่ใส่รหัส (ดังภาพที่ 4.1)

ข้อมูลส่วนบุคคล	สำหรับเจ้าหน้าที่
1. เพศ () ชาย () หญิง	เพศ ☐
2. อายุ () 20-25 ปี () 26-30 ปี () 31-35 ปี	อายุ ☐
3. ระดับการศึกษาสูงสุด () มัธยมศึกษา () ปริญญาตรี () ปริญญาโท () ปริญญาเอก	การศึกษา ☐

ภาพที่ 4.1 ตัวอย่างบางส่วนของแบบสอบถามที่มีการกำหนดช่องสำหรับให้เจ้าหน้าที่ใส่รหัส

ในการกำหนดรหัสให้ข้อมูลเชิงคุณภาพ ตัวแปรที่เป็นข้อความเมื่อแปลงรหัสเป็นตัวเลข จำนวนหลักของตัวเลขควรเท่ากับจำนวนทางเลือกของคำตอบ จากภาพที่ 4.1 ตัวแปรเพศ จะมีค่ารหัส 2 ตัวเท่ากับจำนวนคำตอบ เช่น 1 = เพศชาย และ 2 = เพศหญิง ซึ่งทั้งสองคำตอบเป็นเลข 1 หลัก จึงใช้ช่องสี่เหลี่ยม 1 ช่องสำหรับให้เจ้าหน้าที่ใส่รหัสคำตอบ

การกำหนดรหัสตัวแปรให้ข้อมูลจะขึ้นอยู่กับชนิดของคำถามในแบบสอบถาม การกำหนดรหัสโดยแบ่งตามชนิดของคำถามที่มักถูกใช้บ่อยๆ แบ่งได้ดังนี้

(1) คำถามปลายปิดที่ให้เลือกตอบเพียงข้อเดียว

คำถามประเภทนี้ผู้ตอบแบบสอบถามเลือกคำตอบได้เพียงข้อเดียว เช่น แบบสอบถามเกี่ยวกับข้อมูลส่วนบุคคลที่ให้ผู้ตอบแบบสอบถามตอบว่าเป็นเพศใด (ชาย, หญิง) มีอาชีพอะไร (เช่น นิสิตนักศึกษา, อาจารย์, เจ้าหน้าที่ และ ลูกจ้าง) จากภาพที่ 4.1 สามารถกำหนดรหัสให้แต่ละตัวแปรได้ ดังนี้

ตัวแปรเพศ

- 1 = เพศชาย
2 = เพศหญิง

ตัวแปรอายุ

- 1 = อายุ 20-25 ปี
2 = อายุ 26-30 ปี
3 = อายุ 31-35 ปี

ตัวแปรระดับการศึกษา

- 1 = มัธยมศึกษา
2 = ปริญญาตรี
3 = ปริญญาโท
4 = ปริญญาเอก

(2) คำถามปลายปิดที่ให้เลือกตอบได้หลายข้อ

คำถามประเภทนี้ผู้ตอบแบบสอบถามสามารถเลือกคำตอบได้มากกว่า 1 ข้อ เช่น แบบสอบถามเกี่ยวกับพฤติกรรมกรรมการบริโภค

ท่านเคยดื่มเครื่องดื่มชูกำลังยี่ห้อใดบ้าง (เลือกตอบได้มากกว่า 1 ข้อ)

- ☐ กระทิงแดง ☐ M-150 ☐ คาราบาวแดง ☐ ลิโพวิตัน-ดี

พิจารณาคำถามข้อนี้จะพบว่ามีคำตอบให้เลือก 4 ข้อจึงสร้างตัวแปรได้ 4 ตัวแปร โดยที่แต่ละตัวแปรจะมีค่าได้ 2 ค่าเท่านั้น คือ 0 = ไม่เคยดื่ม และ 1 = เคยดื่ม ดังนี้

ตัวแปรกระทิงแดง

- 0 = ไม่เคยดื่มกระทิงแดง
1 = เคยดื่มกระทิงแดง

ตัวแปรM-150

- 0 = ไม่เคยดื่ม M-150
1 = เคยดื่ม M-150

ตัวแปรคาราบาวแดง

- 0 = ไม่เคยดื่มคาราบาวแดง
1 = เคยดื่มคาราบาวแดง

ตัวแปรลิโพวิตัน-ดี

- 0 = ไม่เคยดื่มลิโพวิตัน-ดี
1 = เคยดื่มลิโพวิตัน-ดี

(3) คำถามที่ให้ใส่ลำดับโดยให้จัดลำดับครบทุกทางเลือก

คำถามประเภทนี้ให้ผู้ตอบแบบสอบถามจัดลำดับให้ครบทุกทางเลือกที่ให้มีมา เช่น เรียงลำดับความชอบ เรียงลำดับความสำคัญ โดย กำหนดให้มีจำนวนตัวแปร = จำนวนทางเลือก ยกตัวอย่างเช่น

กรุณาเรียงลำดับความสำคัญปัจจัยที่มีผลต่อการเลือกซื้อไอศกรีมของท่าน

ให้ลำดับที่ 1 หมายถึง สำคัญมากที่สุด และลำดับที่ 3 หมายถึง สำคัญน้อยที่สุด

- ☐ รสชาติ ☐ เนื้อสัมผัส ☐ ราคา

วิธีการกำหนดตัวแปรทำได้ 2 แบบ คือ

แบบที่ 1 กำหนดค่ารหัสตัวแปรเป็นสเกลนามกำหนด (Nominal scale)

แบบที่ 2 กำหนดค่ารหัสตัวแปรเป็นสเกลอันดับ (Ordinal scale)

แบบที่ 1 กำหนดค่ารหัสตัวแปรเป็นสเกลนามกำหนด (Nominal scale) จากตัวอย่างที่ให้ผู้ตอบแบบสอบถามเรียงลำดับความสำคัญของปัจจัยที่มีผลต่อการเลือกซื้อไอศกรีม มี 3 ตัวเลือก คือ รสชาติ เนื้อสัมผัส และราคา ในกรณีนี้จะกำหนดตัวแปร จำนวน 3 ตัวแปร ตามลำดับความสำคัญ ได้แก่ ตัวแปรลำดับที่ 1, ตัวแปรลำดับที่ 2 และ ตัวแปรลำดับที่ 3 และกำหนดค่ารหัสของแต่ละตัวแปรเป็นสเกลนามกำหนด (1 = รสชาติ, 2 = เนื้อสัมผัส, 3 = ราคา) ดังนี้

ตัวแปร ลำดับที่ 1

- 1 = เลือกรสชาติเป็นลำดับที่ 1
- 2 = เลือกเนื้อสัมผัสเป็นลำดับที่ 1
- 3 = เลือกราคาเป็นลำดับที่ 1

ตัวแปร ลำดับที่ 2

- 1 = เลือกรสชาติเป็นลำดับที่ 2
- 2 = เลือกเนื้อสัมผัสเป็นลำดับที่ 2
- 3 = เลือกราคาเป็นลำดับที่ 2

ตัวแปร ลำดับที่ 3

- 1 = เลือกรสชาติเป็นลำดับที่ 3
- 2 = เลือกเนื้อสัมผัสเป็นลำดับที่ 3
- 3 = เลือกราคาเป็นลำดับที่ 3

แบบที่ 2 กำหนดค่ารหัสตัวแปรเป็นสเกลอันดับ (Ordinal scale) จากตัวอย่างที่ให้ผู้ตอบแบบสอบถามเรียงลำดับความสำคัญของปัจจัยที่มีผลต่อการเลือกซื้อไอศกรีม มี 3 ตัวเลือก คือ รสชาติ เนื้อสัมผัส และราคา ในกรณีนี้จะกำหนดตัวแปร จำนวน 3 ตัวแปร ตามจำนวนทางเลือกคำตอบ ได้แก่ ตัวแปรรสชาติ, ตัวแปรเนื้อสัมผัส และ ตัวแปรราคา และกำหนดค่ารหัสของแต่ละตัวแปรเป็นสเกลอันดับ (1 = สำคัญมากที่สุด, 2 = สำคัญปานกลาง, 3 = สำคัญน้อยที่สุด) ดังนี้

ตัวแปรรสชาติ

- 1 = สำคัญมากที่สุด
- 2 = สำคัญปานกลาง
- 3 = สำคัญน้อยที่สุด

ตัวแปรเนื้อสัมผัส

- 1 = สำคัญมากที่สุด
- 2 = สำคัญปานกลาง
- 3 = สำคัญน้อยที่สุด

ตัวแปรราคา

- 1 = สำคัญมากที่สุด
- 2 = สำคัญปานกลาง
- 3 = สำคัญน้อยที่สุด

(4) คำถามที่ให้ใส่ลำดับที่โดยให้จัดลำดับน้อยกว่าจำนวนทางเลือก

คำถามประเภทนี้ให้ผู้ตอบแบบสอบถามจัดลำดับน้อยกว่าทางเลือกที่ให้มา เช่น มีคำตอบให้เลือก 5 คำตอบแต่ให้ผู้ตอบแบบสอบถามเรียงลำดับ 3 ลำดับแรก ยกตัวอย่างเช่น

กรุณาเรียงลำดับรสชาติน้ำอัดลมที่ท่านชอบ 3 ลำดับแรก ให้ลำดับที่ 1 หมายถึงชอบมากที่สุด

☐ โคลา ☐ น้ำส้ม ☐ น้ำแดง ☐ น้ำเขียว ☐ น้ำองุ่น

วิธีการกำหนดตัวแปรทำได้ 2 แบบ คือ

แบบที่ 1 กำหนดให้มีจำนวนตัวแปร = จำนวนทางเลือก

แบบที่ 2 กำหนดให้มีจำนวนตัวแปร = จำนวนลำดับที่ให้จัด

แบบที่ 1 กำหนดให้มีจำนวนตัวแปร = จำนวนทางเลือก จากตัวอย่างที่ให้ผู้ตอบแบบสอบถามเรียงลำดับรสชาติน้ำอัดลมที่ชอบ 3 ลำดับแรก จาก 5 ตัวเลือก (โคลา, น้ำส้ม, น้ำแดง, น้ำเขียว, น้ำอู่น) ในกรณีนี้จะกำหนดตัวแปร จำนวน 5 ตัวแปร **ตามจำนวนทางเลือกคำตอบ** ได้แก่ ตัวแปรโคลา, ตัวแปรน้ำส้ม, ตัวแปรน้ำแดง, ตัวแปรน้ำเขียว และตัวแปรน้ำอู่น และกำหนดค่ารหัสของแต่ละตัวแปรเป็นสเกลอันดับ (1=ชอบมากเป็นลำดับที่ 1, 2=ชอบมากเป็นลำดับที่ 2, 3=ชอบมากเป็นลำดับที่ 3 และ 0=ผู้ตอบแบบสอบถามไม่เลือก) ดังนี้

ตัวแปร โคลา 1 = ชอบมากเป็นลำดับที่ 1 2 = ชอบมากเป็นลำดับที่ 2 3 = ชอบมากเป็นลำดับที่ 3 0 = ผู้ตอบแบบสอบถามไม่เลือก	ตัวแปร น้ำส้ม 1 = ชอบมากเป็นลำดับที่ 1 2 = ชอบมากเป็นลำดับที่ 2 3 = ชอบมากเป็นลำดับที่ 3 0 = ผู้ตอบแบบสอบถามไม่เลือก	ตัวแปร น้ำแดง 1 = ชอบมากเป็นลำดับที่ 1 2 = ชอบมากเป็นลำดับที่ 2 3 = ชอบมากเป็นลำดับที่ 3 0 = ผู้ตอบแบบสอบถามไม่เลือก
ตัวแปร น้ำเขียว 1 = ชอบมากเป็นลำดับที่ 1 2 = ชอบมากเป็นลำดับที่ 2 3 = ชอบมากเป็นลำดับที่ 3 0 = ผู้ตอบแบบสอบถามไม่เลือก	ตัวแปร น้ำอู่น 1 = ชอบมากเป็นลำดับที่ 1 2 = ชอบมากเป็นลำดับที่ 2 3 = ชอบมากเป็นลำดับที่ 3 0 = ผู้ตอบแบบสอบถามไม่เลือก	

แบบที่ 2 กำหนดให้มีจำนวนตัวแปร = จำนวนลำดับที่ให้จัด จากตัวอย่างที่ให้ผู้ตอบแบบสอบถามเรียงลำดับรสชาติน้ำอัดลมที่ชอบ 3 ลำดับแรก จาก 5 ตัวเลือก (โคลา, น้ำส้ม, น้ำแดง, น้ำเขียว, น้ำอู่น) ในกรณีนี้จะกำหนดตัวแปร จำนวน 3 ตัวแปร **ตามลำดับความสำคัญ** ได้แก่ ตัวแปรลำดับที่ 1, ตัวแปรลำดับที่ 2 และ ตัวแปรลำดับที่ 3 และกำหนดค่ารหัสของแต่ละตัวแปรเป็นสเกลนามกำหนด (1=โคลา, 2=น้ำส้ม, 3=น้ำแดง, 4=น้ำเขียว, 5=น้ำอู่น) ดังนี้

ตัวแปรลำดับที่ 1 1 = โคลา 2 = น้ำส้ม 3 = น้ำแดง 4 = น้ำเขียว 5 = น้ำอู่น	ตัวแปรลำดับที่ 2 1 = โคลา 2 = น้ำส้ม 3 = น้ำแดง 4 = น้ำเขียว 5 = น้ำอู่น	ตัวแปรลำดับที่ 3 1 = โคลา 2 = น้ำส้ม 3 = น้ำแดง 4 = น้ำเขียว 5 = น้ำอู่น
--	--	--

4.3.2 การใช้คำสั่ง SPSS วิเคราะห์สถิติเชิงพรรณนาสำหรับข้อมูลสเกลนามกำหนด

ข้อมูลสเกลนามกำหนดที่มักพบในแบบสอบถาม เช่น เพศ ระดับการศึกษา อาชีพ ความคิดเห็น ซึ่งสถิติเชิงพรรณนาที่ใช้วิเคราะห์ คือ การหาจำนวนหรือความถี่ และร้อยละ โดยคำสั่งในโปรแกรม SPSS คือ

Analyze ➤ Descriptive Statistics ➤ Frequencies...

ตัวอย่าง นักวิจัยทำการสำรวจพฤติกรรมการดูหนังของผู้ตอบแบบสอบถาม 100 คน โดยคำถามในส่วนของข้อมูลส่วนบุคคลที่ใช้สเกลแบบนามกำหนด คือ เพศ ระดับการศึกษา และ อาชีพ ลักษณะคำถามเป็นดังนี้

1. เพศ () ชาย () หญิง
2. ระดับการศึกษา () มัธยมศึกษา () ปริญญาตรี () ปริญญาโท () ปริญญาเอก
3. อาชีพ () นักเรียน () นิสิตนักศึกษา () รับราชการ () รัฐวิสาหกิจ () พนักงานเอกชน

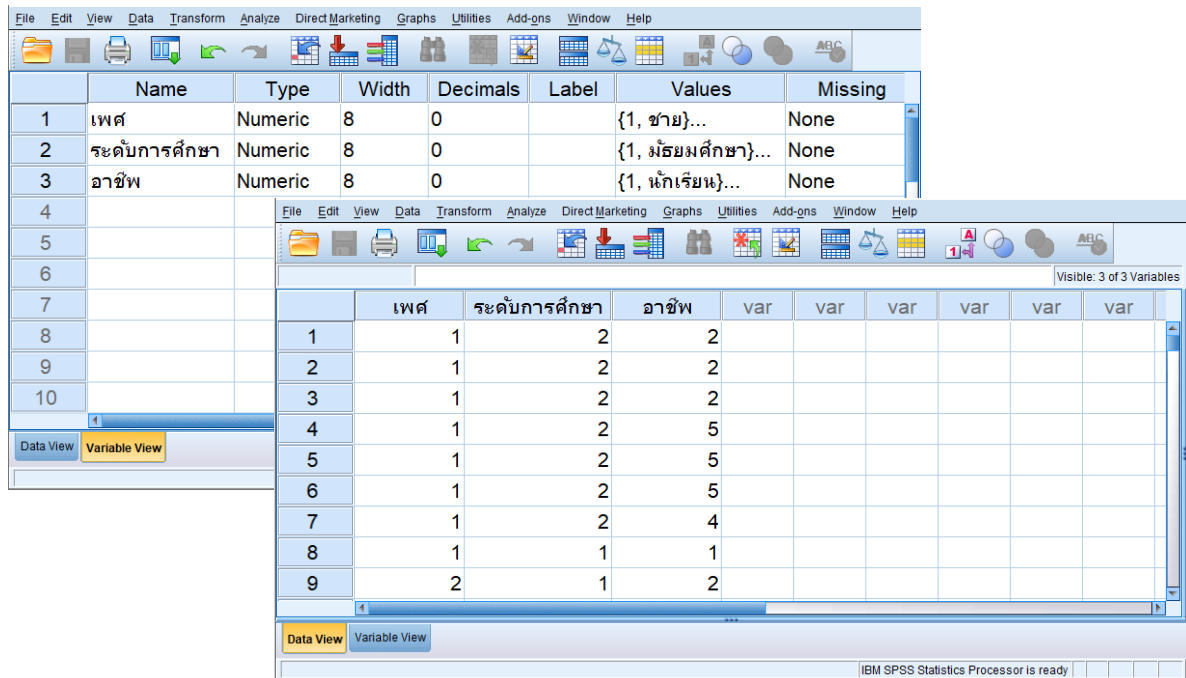
เนื่องจากข้อมูลตัวแปร เพศ ระดับการศึกษา และอาชีพ ทั้งสามตัวแปรเป็นสเกลแบ่งกลุ่ม ดังนั้นสถิติเชิงพรรณนาที่ใช้วิเคราะห์ คือ จำนวนหรือความถี่ และ ร้อยละ

ขั้นตอนการวิเคราะห์จำนวนหรือความถี่ และร้อยละ ด้วยโปรแกรม SPSS

(1) ป้อนข้อมูลของผู้ตอบแบบสอบถามแต่ละคนลงในโปรแกรม SPSS โดยกำหนดค่าให้ตัวแปร var ดังนี้

ข้อมูลจากแบบสอบถาม	ชื่อตัวแปร var	กำหนดค่า Values
เพศ	เพศ	1 = เพศชาย 2 = เพศหญิง
ระดับการศึกษา	ระดับการศึกษา	1 = มัธยมศึกษา 2 = ปริญญาตรี 3 = ปริญญาโท 4 = ปริญญาเอก
อาชีพ	อาชีพ	1 = นักเรียน 2 = นิสิตนักศึกษา 3 = รับราชการ 4 = รัฐวิสาหกิจ 5 = พนักงานเอกชน

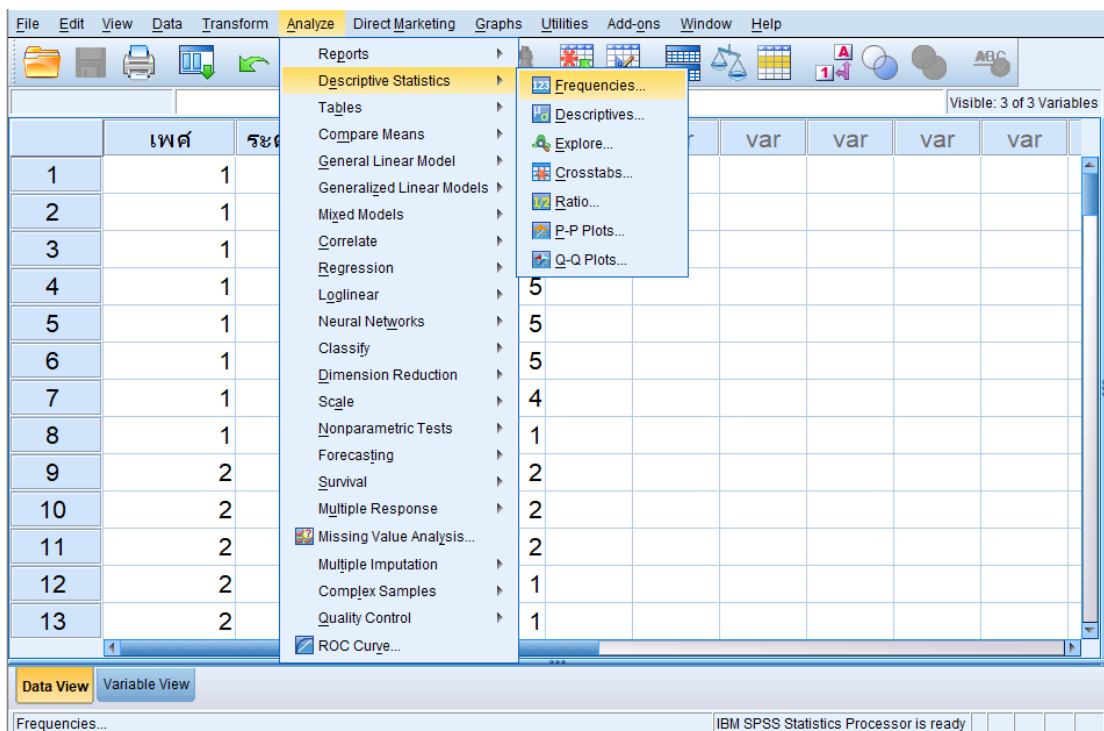
โปรแกรม SPSS for Windows version 19 สามารถพิมพ์ภาษาไทยได้ เมื่อป้อนข้อมูลตัวแปรในหน้าต่าง Variable View และข้อมูลของผู้ตอบแบบสอบถามในหน้าต่าง Data View เสร็จแล้วจะได้ดังภาพที่ 4.2



ภาพที่ 4.2 หน้าต่าง Variable View และ Data View

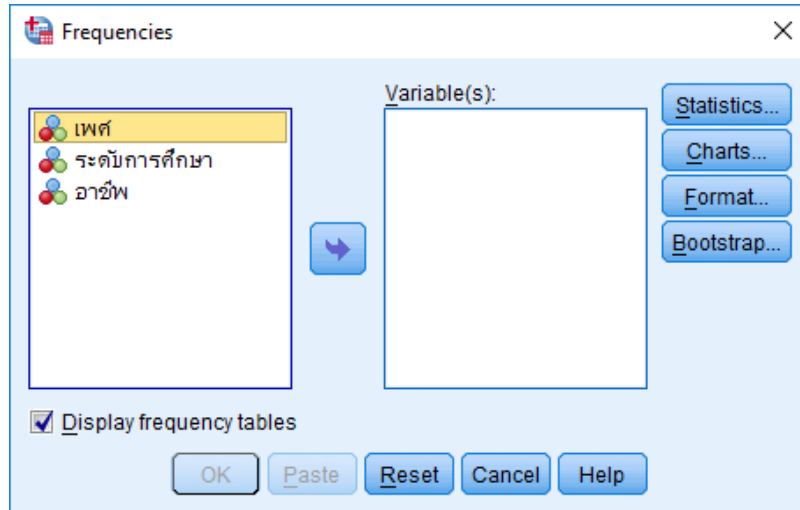
(2) เมื่อป้อนข้อมูลในโปรแกรม SPSS และตรวจสอบความถูกต้องแล้ว ให้ใช้คำสั่ง

Analyze ➤ Descriptive Statistics ➤ Frequencies... ดังภาพที่ 4.3



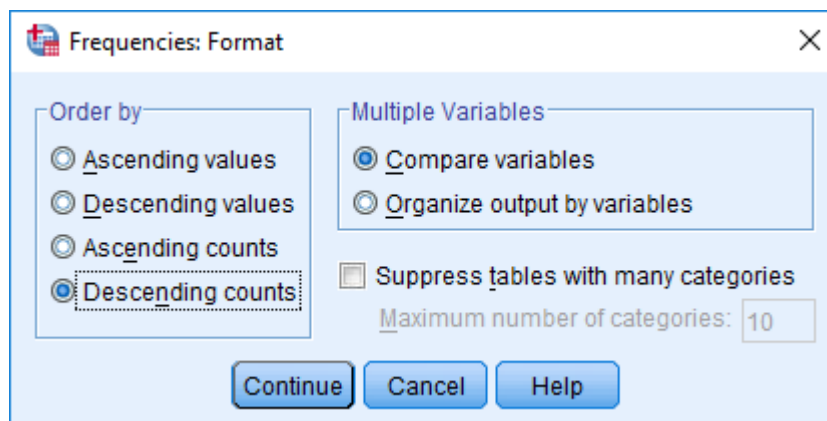
ภาพที่ 4.3 เมนูคำสั่ง Analyze ➤ Descriptive Statistics ➤ Frequencies...

(3) ในหน้าจอคำสั่ง Frequencies (ภาพที่ 4.4) ให้เลือกตัวแปรเพศ ระดับการศึกษา อาชีพ จากกล่อง ด้านซ้ายมือไปใส่ในกล่อง Variable(s) ที่อยู่ด้านขวามือ โดยใช้ลูกศรตัวกลาง



ภาพที่ 4.4 หน้าจอคำสั่ง Frequencies

(4) กดคลิกเลือก **Format...** เพื่อกำหนดรูปแบบตารางที่ต้องการนำเสนอผลการวิเคราะห์ จะปรากฏ หน้าจอดังภาพที่ 4.5 กดคลิกเลือก Descending counts แล้วกด Continue และ OK จะได้ผลลัพธ์ดังภาพที่ 4.6



ภาพที่ 4.5 หน้าจอ Frequencies: Format

ความหมายของ Order by ทั้ง 4 ชนิด คือ Ascending values เรียงตามค่ารหัสตัวแปรจากน้อยไปมาก, Descending values เรียงตามค่ารหัสตัวแปรจากมากไปน้อย, Ascending counts เรียงตามความถี่จากน้อยไปมาก และ Descending counts เรียงตามความถี่จากมากไปน้อย

Frequencies

Statistics

		เพศ	ระดับการศึกษา	อาชีพ
N	Valid	100	100	100
	Missing	0	0	0

Frequency Table

เพศ

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	ชาย	53	53.0	53.0	53.0
	หญิง	47	47.0	47.0	100.0
	Total	100	100.0	100.0	

ระดับการศึกษา

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	ปริญญาตรี	44	44.0	44.0	44.0
	มัธยมศึกษา	31	31.0	31.0	75.0
	ปริญญาโท	16	16.0	16.0	91.0
	ปริญญาเอก	9	9.0	9.0	100.0
	Total	100	100.0	100.0	

อาชีพ

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	นิสิตนักศึกษา	28	28.0	28.0	28.0
	นักเรียน	25	25.0	25.0	53.0
	พนักงานบริษัทเอกชน	19	19.0	19.0	72.0
	รัฐวิสาหกิจ	15	15.0	15.0	87.0
	รับราชการ	13	13.0	13.0	100.0
	Total	100	100.0	100.0	

ภาพที่ 4.6 ผลลัพธ์ในรูปแบบตารางจากการวิเคราะห์ความถี่โดยโปรแกรม SPSS

เมื่อพิจารณาจากภาพที่ 4.6 จะพบว่ากลุ่มตัวอย่างเป็นเพศชายร้อยละ 53 และ เพศหญิงร้อยละ 47, กลุ่มตัวอย่างส่วนใหญ่ร้อยละ 44 มีการศึกษาระดับปริญญาตรี และร้อยละ 28 เป็นนิสิตและนักศึกษา (ข้อสังเกต คือ ค่าร้อยละ (Percent) และค่า Valid Percent เท่ากัน เพราะผู้ตอบแบบสอบถามทุกคนตอบคำถามข้อนี้)

(5) หากต้องการนำเสนอผลการวิเคราะห์ในรูปแบบกราฟแท่ง (Bar charts) ให้ใช้คำสั่ง

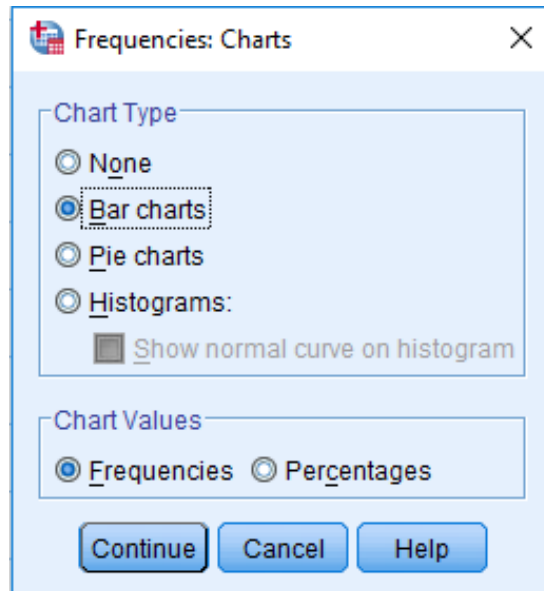
Analyze ➤ Descriptive Statistics ➤ Frequencies...

คลิกเลือก

Charts...

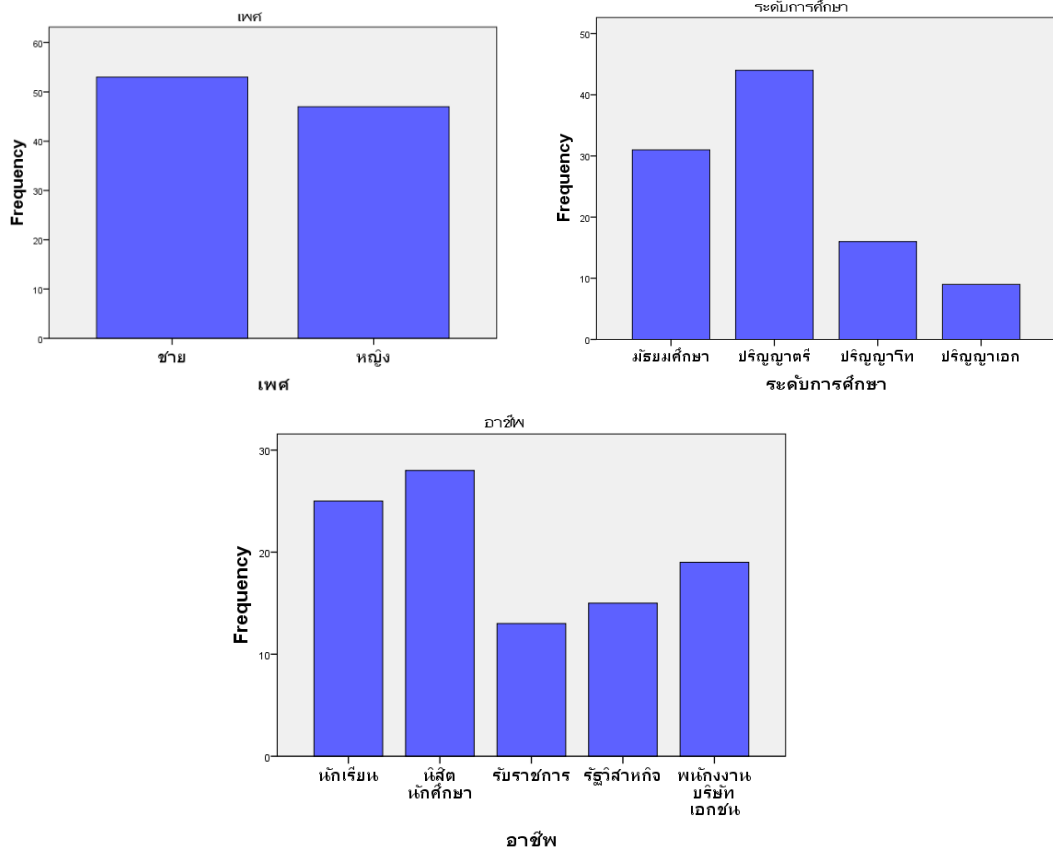
จะได้หน้าจอตั้ง

ภาพที่ 4.7 กดเลือก Bar charts, กด Continue และกด OK จะได้ผลลัพธ์ดังภาพที่ 4.8 (ถ้าผู้ใช้เลือก Bar charts หรือ Pie charts จะสามารถเลือกให้แสดงข้อมูลเป็น ความถี่ หรือ เปอร์เซ็นต์ได้)



ภาพที่ 4.7 หน้าจอ Frequencies: Charts

Bar Chart



ภาพที่ 4.8 ผลลัพธ์ในรูปแบบแผนภูมิแท่งจากการวิเคราะห์ความถี่โดยโปรแกรม SPSS

4.3.3 การใช้คำสั่ง SPSS วิเคราะห์สถิติเชิงพรรณนาสำหรับข้อมูลสเกลอันดับ

ตัวอย่างสเกลอันดับที่มักใช้ในการสร้างแบบสอบถาม เช่น อายุที่แบ่งเป็นช่วง (เช่น 15-20 ปี, 21-25 ปี และ 26-30 ปี) รายได้ต่อเดือน (เช่น น้อยกว่า 5,000 บาท, 5,000-10,000 บาท และมากกว่า 10,000 บาท ขึ้นไป) และ ระดับการศึกษา (เช่น มัธยมศึกษา, ปริญญาตรี, ปริญญาโท และ ปริญญาเอก) ทั้งนี้จะเห็นว่าอายุที่เป็นช่วง, รายได้ต่อเดือน และ ระดับการศึกษา นักวิจัยอาจพิจารณาว่าเป็นได้ทั้งข้อมูลสเกลนามกำหนดและสเกลอันดับ เนื่องจากเราสามารถแบ่งกลุ่มแล้วเรียงลำดับอายุและเงินเดือนจากน้อยไปมาก และเรียงลำดับระดับการศึกษาจากต่ำไปสูงได้

ข้อมูลสเกลอันดับนี้จะใช้วิธีการวิเคราะห์สถิติเชิงพรรณนา ได้แก่ จำนวนหรือความถี่ และ ร้อยละ เช่นเดียวกับสเกลนามกำหนด โดยคำสั่งในโปรแกรม SPSS คือ

Analyze ➤ Descriptive Statistics ➤ Frequencies...

ตัวอย่าง จากข้อมูลในข้อ 4.3.2 ถ้าทำการสำรวจพฤติกรรมการดูหนังของผู้ตอบแบบสอบถาม 100 คน โดยการตั้งคำถามในส่วนของคุณลักษณะส่วนบุคคลด้าน รายได้ต่อเดือน ด้วยสเกลอันดับ ลักษณะคำถามเป็นดังนี้

ท่านมีรายได้ต่อเดือนเท่าไร

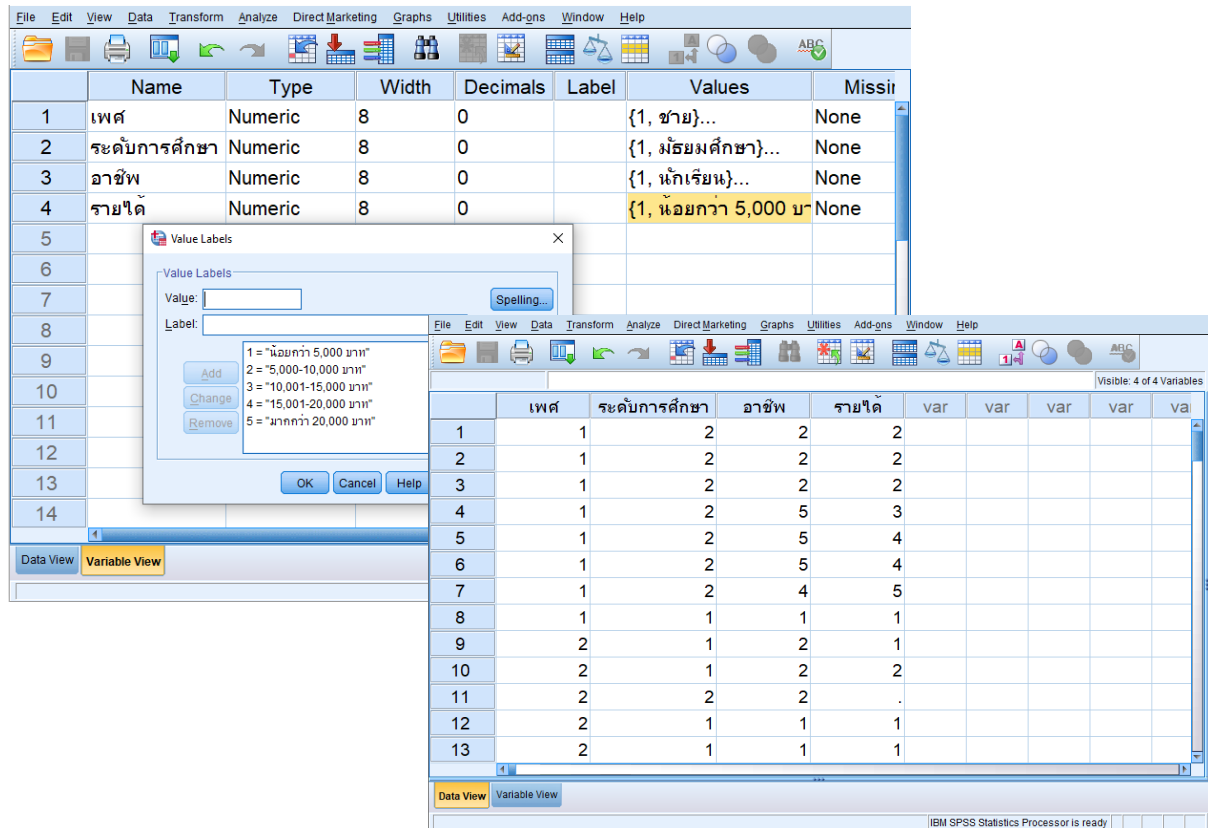
- () น้อยกว่า 5,000 บาท () 5,000 - 10,000 บาท () 10,001 - 15,000 บาท
() 15,001 - 20,000 บาท () มากกว่า 20,000 บาท

ขั้นตอนการวิเคราะห์จำนวนหรือความถี่ และร้อยละ ด้วยโปรแกรม SPSS

(1) ป้อนข้อมูลรายได้ของผู้ตอบแบบสอบถามแต่ละคนในโปรแกรม SPSS โดยทำการเพิ่มตัวแปร var กำหนดให้ชื่อ รายได้ และ กำหนดค่าให้รหัสตัวแปรรายได้ ดังนี้

ข้อมูลจากแบบสอบถาม	ชื่อตัวแปร var	กำหนดค่า Values
รายได้	รายได้	1 = น้อยกว่า 5,000 บาท 2 = 5,000-10,000 บาท 3 = 10,001-15,000 บาท 4 = 15,001-20,000 บาท 5 = มากกว่า 20,000 บาท

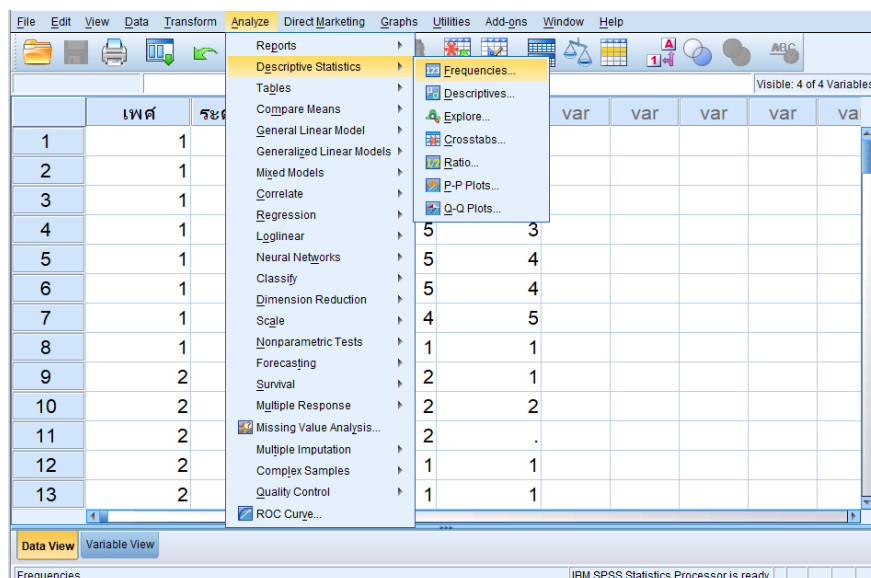
เมื่อกำหนดตัวแปรรายได้ในหน้าต่าง Variable View และป้อนข้อมูลรายได้ของผู้ตอบแบบสอบถามในหน้าต่าง Data View เสร็จแล้วจะได้ดังภาพที่ 4.9



ภาพที่ 4.9 หน้าต่าง Variable View และ Data View เมื่อป้อนข้อมูลรายได้ของผู้ตอบแบบสอบถาม

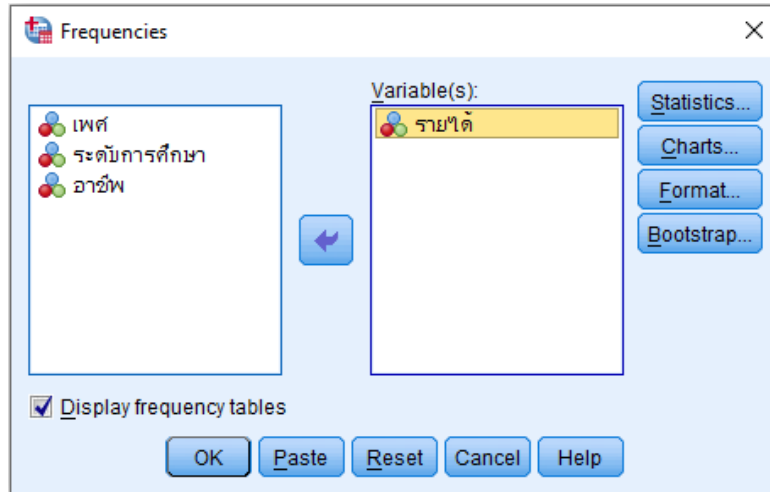
(2) เมื่อป้อนข้อมูลรายได้ในโปรแกรม SPSS และตรวจสอบความถูกต้องแล้ว ให้ใช้คำสั่ง

Analyze ➤ Descriptive Statistics ➤ Frequencies... ดังภาพที่ 4.10



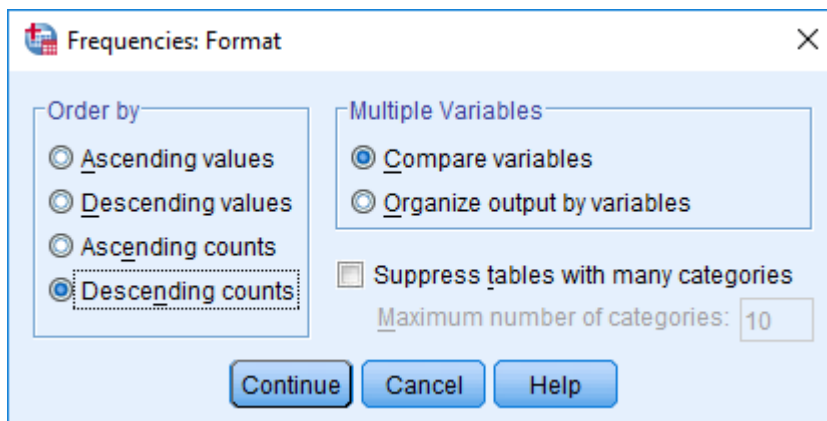
ภาพที่ 4.10 เมนูคำสั่ง Analyze ➤ Descriptive Statistics ➤ Frequencies...

(3) ในหน้าจอคำสั่ง Frequencies (ภาพที่ 4.11) ให้เลือกตัวแปรรายได้ที่อยู่ในกล่องด้านซ้ายมือไปใส่ในกล่อง Variable(s) ที่อยู่ด้านขวามือ (ในกรณีที่ต้องการคำนวณความถี่และร้อยละของตัวแปร เพศ ระดับการศึกษา อาชีพ และรายได้ ผู้วิจัยสามารถวิเคราะห์ได้พร้อมกันเลย โดยเลือกตัวแปรทั้งหมดจากกล่องด้านซ้ายมือไปใส่ในกล่องด้านขวามือ)



ภาพที่ 4.11 หน้าจอคำสั่ง Frequencies

(4) ให้กดคลิกเลือก **Format...** เพื่อกำหนดรูปแบบตารางที่ต้องการนำเสนอผลการวิเคราะห์ คลิกเลือก Descending counts ดังภาพที่ 4.12 แล้วกด Continue และ OK จะได้ผลลัพธ์ดังภาพที่ 4.13



ภาพที่ 4.12 หน้าจอ Frequencies: Format

Frequencies

Statistics

รายได้

N	Valid	99
	Missing	1

รายได้

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	5,000-10,000 บาท	34	34.0	34.3	34.3
	15,001-20,000 บาท	19	19.0	19.2	53.5
	น้อยกว่า 5,000 บาท	17	17.0	17.2	70.7
	10,001-15,000 บาท	16	16.0	16.2	86.9
	มากกว่า 20,000 บาท	13	13.0	13.1	100.0
	Total	99	99.0	100.0	
Missing	System	1	1.0		
Total		100	100.0		

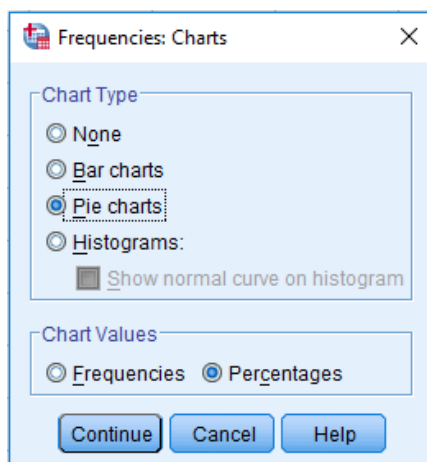
ภาพที่ 4.13 ผลลัพธ์การหาความถี่ตัวแปรรายได้โดยใช้โปรแกรม SPSS

จากภาพที่ 4.13 จะพบว่าผู้ทดสอบ 100 คน มี 99 คนที่ตอบคำถามเกี่ยวกับรายได้ และมี 1 คนไม่ได้ตอบคำถามนี้ (จากภาพที่ 4.9 คือ ผู้ตอบแบบสอบถามคนที่ 11) ถือเป็นข้อมูล Missing data ในคอลัมน์ Percent จะเป็นการแสดงค่าร้อยละโดยคำนวณจากผู้ตอบแบบสอบถาม 100 คน ในขณะที่คอลัมน์ Valid Percent จะเป็นการแสดงค่าร้อยละโดยคำนวณจากผู้ตอบแบบสอบถาม 99 คน

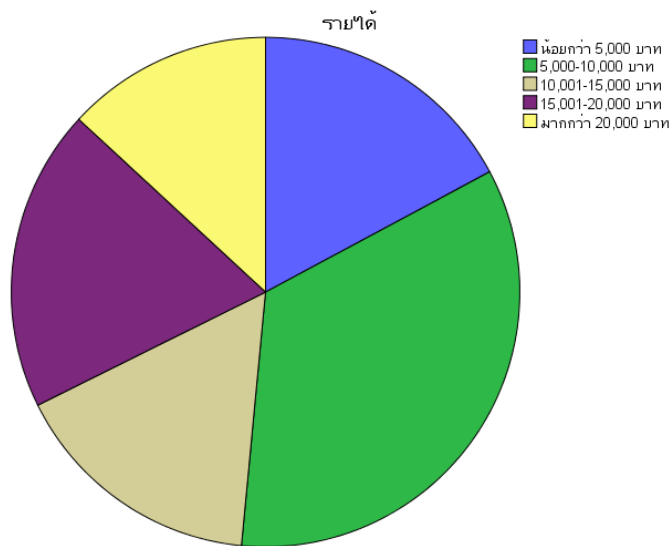
(5) หากต้องการนำเสนอผลการวิเคราะห์ในรูปแบบแผนภูมิวงกลม (Pie chart) ให้ใช้คำสั่ง

Analyze ► Descriptive Statistics ► Frequencies... คลิกเลือก Charts...

จะได้หน้าจอตั้งภาพที่ 4.14 กดเลือก Pie charts, เลือกค่าแบบ Percentages, กด Continue และกด OK จะได้ผลลัพธ์แสดงดังภาพที่ 4.15

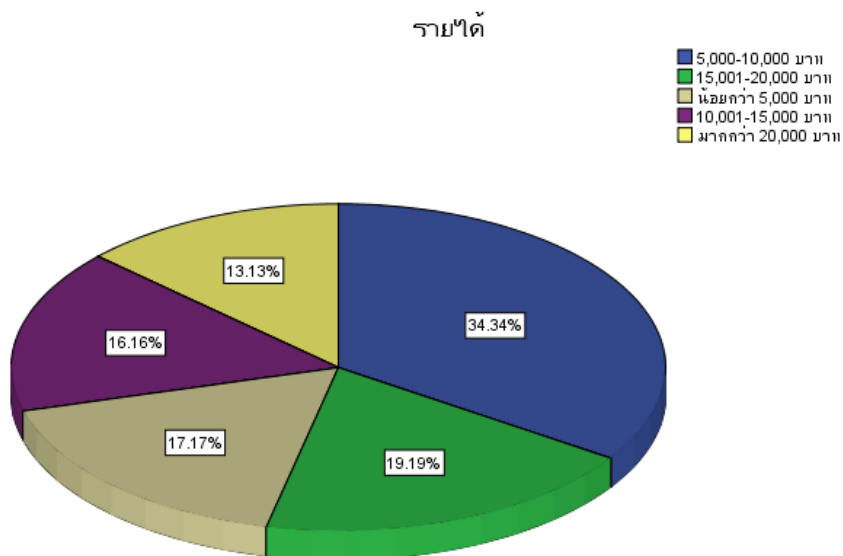


ภาพที่ 4.14 หน้าจอ Frequencies: Charts



ภาพที่ 4.15 ผลลัพธ์ในรูปแบบแผนภูมิวงกลมจากการวิเคราะห์ความถี่โดยโปรแกรม SPSS

ผู้ใช้งานสามารถแก้ไขรูปภาพวงกลมที่ออกมาได้ โดยกดดับเบิลคลิกที่ตัวกราฟ จะปรากฏกล่อง Chart Editor เพื่อให้ผู้ใช้งานสามารถแก้ไขรูปแบบกราฟตามต้องการ เช่น แก้ไขรูปแบบให้เป็นกราฟ 3 มิติ และเพิ่มค่าข้อมูลแสดงในกราฟ ดังภาพที่ 4.16



ภาพที่ 4.16 ผลลัพธ์การแก้ไขแผนภูมิวงกลมให้เป็นกราฟ 3 มิติ และเพิ่มค่าข้อมูลในกราฟ

4.4 สถิติเชิงพรรณนาสำหรับข้อมูลเชิงปริมาณ

ข้อมูลเชิงปริมาณเป็นข้อมูลตัวเลขที่สามารถวัดค่าได้ว่ามีค่ามากหรือน้อย ข้อมูลที่ได้จากสเกลอันดับและสเกลอัตราส่วน คือ ข้อมูลเชิงปริมาณ เช่น รายได้ อายุ เกรดเฉลี่ยสะสม ยอดขาย ปริมาณไขมัน ปริมาณโปรตีน ราคาขาย ระดับความพึงพอใจ และ คะแนนความชอบ ค่าสถิติเชิงพรรณนาที่ใช้วิเคราะห์ข้อมูลเชิงปริมาณ ได้แก่

4.4.1 ค่ากลางของข้อมูล ได้แก่ ค่าเฉลี่ย, ค่ามัธยฐาน และ ค่าฐานนิยม เช่น ถ้ามีข้อมูล 7 ตัว (3 3 4 5 6 7 8) ค่ากลางของข้อมูลแสดงในตารางที่ 4.1

ตารางที่ 4.1 ค่าเฉลี่ย ค่ามัธยฐาน และค่าฐานนิยมของข้อมูล 3 3 4 5 6 7 8

ค่ากลาง	คำอธิบาย	ค่าที่ได้
ค่าเฉลี่ย (Average)	เป็นค่ากลางที่นิยมใช้มากที่สุด ได้จากผลรวมของข้อมูลหารด้วยจำนวนข้อมูล	5.14
ค่ามัธยฐาน (Median)	เป็นค่าของข้อมูลที่มีตำแหน่งตรงกลางของข้อมูลเมื่อนำข้อมูลมาเรียงลำดับจากค่าน้อยไปมาก ดังนั้นจะมีข้อมูลครึ่งหนึ่งน้อยกว่าค่ามัธยฐาน และอีกครึ่งหนึ่งมีค่ามากกว่ามัธยฐาน	5
ค่าฐานนิยม (Mode)	ค่าของข้อมูลที่เกิดขึ้นบ่อยที่สุด หรือมีความถี่สูงสุด	3

4.4.2 ค่าการกระจายของข้อมูลเชิงปริมาณ

ข้อมูลเชิงปริมาณนอกจากจะสรุปลักษณะของข้อมูลด้วยค่ากลางแล้วควรจะต้องสรุปด้วยค่าการกระจายของข้อมูล ดังนี้

- ค่าพิสัย (Range) คำนวณได้จาก ค่าสูงสุด – ค่าต่ำสุด

- ค่าแปรปรวน (Variance) เป็นค่าการวัดการกระจายที่นิยมใช้มาก พิจารณาจากค่าเฉลี่ยของความแตกต่างระหว่างค่าข้อมูลแต่ละค่ากับค่าเฉลี่ยยกกำลังสอง ถ้าค่าแปรปรวนมีค่ามากแสดงว่าข้อมูลชุดนั้นมีการกระจายมาก

$$\text{ค่าแปรปรวนของประชากร} = \sigma^2 = \sum_{i=1}^N \frac{(x_i - \mu)^2}{N}$$

$$\text{ค่าแปรปรวนของตัวอย่าง} = S^2 = \sum_{i=1}^n \frac{(x_i - \bar{x})^2}{n-1}$$

โดยที่ μ = ค่าเฉลี่ยประชากร N = ขนาดประชากร

$\bar{x}$ = ค่าเฉลี่ยตัวอย่าง n = ขนาดตัวอย่าง

- ค่าส่วนเบี่ยงเบนมาตรฐาน (Standard deviation: S.D.) เป็นค่าฐานที่สองของค่าแปรปรวน มีหน่วยเช่นเดียวกับหน่วยของข้อมูล เช่น ปริมาณโปรตีนมีหน่วยเป็นร้อยละ ค่าส่วนเบี่ยงเบนมาตรฐานของปริมาณโปรตีนก็ต้องมีหน่วยเป็นร้อยละเช่นกัน

$$\text{ค่าส่วนเบี่ยงเบนมาตรฐานของประชากร} = \sigma = \sqrt{\sum_{i=1}^N \frac{(x_i - \mu)^2}{N}}$$

$$\text{ค่าส่วนเบี่ยงเบนมาตรฐานของตัวอย่าง} = S = \sqrt{\sum_{i=1}^n \frac{(x_i - \bar{x})^2}{n-1}}$$

ตัวอย่าง จากข้อ 4.3 ที่ได้ยกตัวอย่างการสร้างคำถามในแบบสอบถามด้วยสเกลนามกำหนดและสเกลอันดับ เพื่อทำการสำรวจพฤติกรรมการดูหนังของผู้ตอบแบบสอบถามจำนวน 100 คนนั้น นักวิจัยสามารถสร้างคำถามที่ใช้สเกลอัตราส่วนได้ เช่น การถามอายุของผู้ตอบแบบสอบถามและการสอบถามความถี่ในการไปดูหนังโดยใช้คำถามปลายเปิดเพื่อให้ผู้ตอบแบบสอบถามระบุเป็นตัวเลขเอง ดังนี้

ข้อมูลส่วนบุคคล

1. เพศ () ชาย () หญิง
2. อายุ ปี
3. ระดับการศึกษา () มัธยมศึกษา () ปริญญาตรี () ปริญญาโท () ปริญญาเอก
4. อาชีพ () นักเรียน () นิสิตนักศึกษา () รับราชการ () รัฐวิสาหกิจ () พนักงานเอกชน

ข้อมูลพฤติกรรมการดูหนัง

กรุณาระบุความถี่ในการไปดูหนังของท่าน จำนวน.....ครั้งต่อเดือน

4.4.3 ขั้นตอนการวิเคราะห์หาค่ากลางและค่าการกระจายด้วยโปรแกรม SPSS

ทำได้ 2 วิธีโดยใช้คำสั่ง **Frequencies** และคำสั่ง **Descriptives** มีรายละเอียด ดังนี้

4.4.3.1 การวิเคราะห์หาค่ากลางและค่าการกระจายด้วยคำสั่ง Frequencies

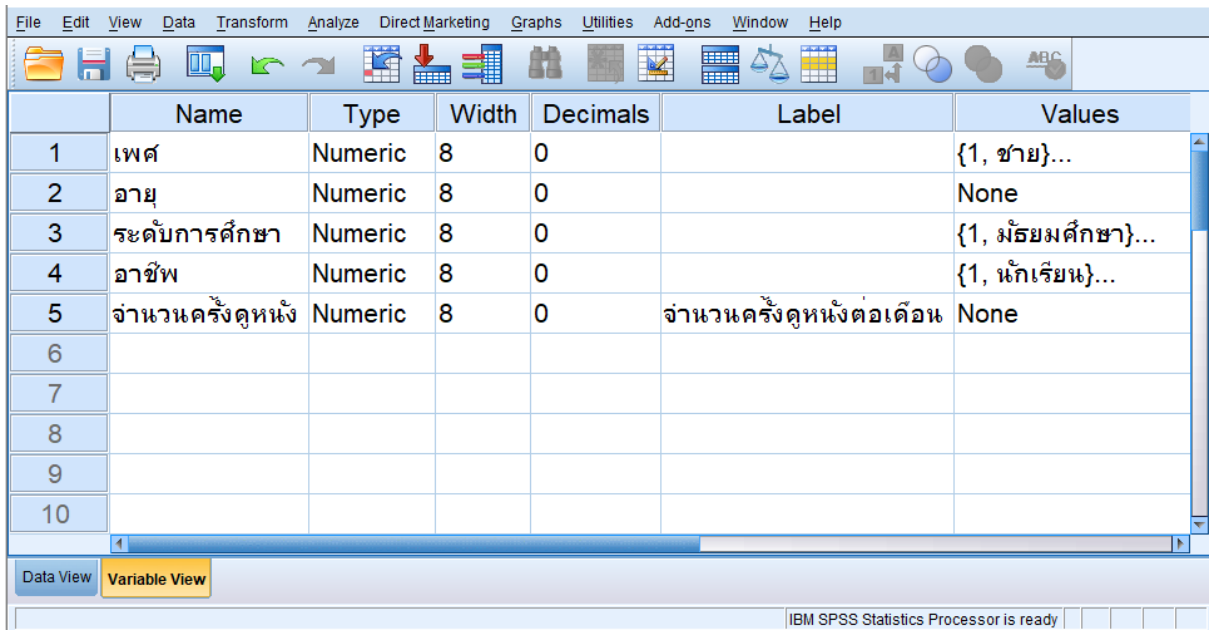
คำสั่ง **Frequencies** นอกจากจะสามารถหาความถี่หรือจำนวน และร้อยละ ดังที่ได้อธิบายในข้อ 4.3 แล้ว ยังสามารถใช้หาค่ากลางและค่าการกระจายของข้อมูลเชิงปริมาณได้ด้วย

ขั้นตอนการวิเคราะห์หาค่ากลางและค่าการกระจายด้วยคำสั่ง Frequencies

(1) ป้อนข้อมูล เพศ อายุ ระดับการศึกษา อาชีพ และจำนวนครั้งในการดูหนังต่อเดือนของผู้ตอบแบบสอบถามแต่ละคน ในโปรแกรม SPSS โดยทำการกำหนดชื่อตัวแปร var และค่า Values ดังนี้

ข้อมูลจากแบบสอบถาม	ชื่อตัวแปร var	กำหนดค่า Values
เพศ	เพศ	1 = เพศชาย 2 = เพศหญิง
อายุ	อายุ	ไม่ต้องระบุ
ระดับการศึกษา	ระดับการศึกษา	1 = มัธยมศึกษา 2 = ปริญญาตรี 3 = ปริญญาโท 4 = ปริญญาเอก
อาชีพ	อาชีพ	1 = นักเรียน 2 = นิสิตนักศึกษา 3 = รับราชการ 4 = รัฐวิสาหกิจ 5 = พนักงานเอกชน
จำนวนครั้งในการดูหนังต่อเดือน	จำนวนครั้งดูหนังต่อเดือน	ไม่ต้องระบุ

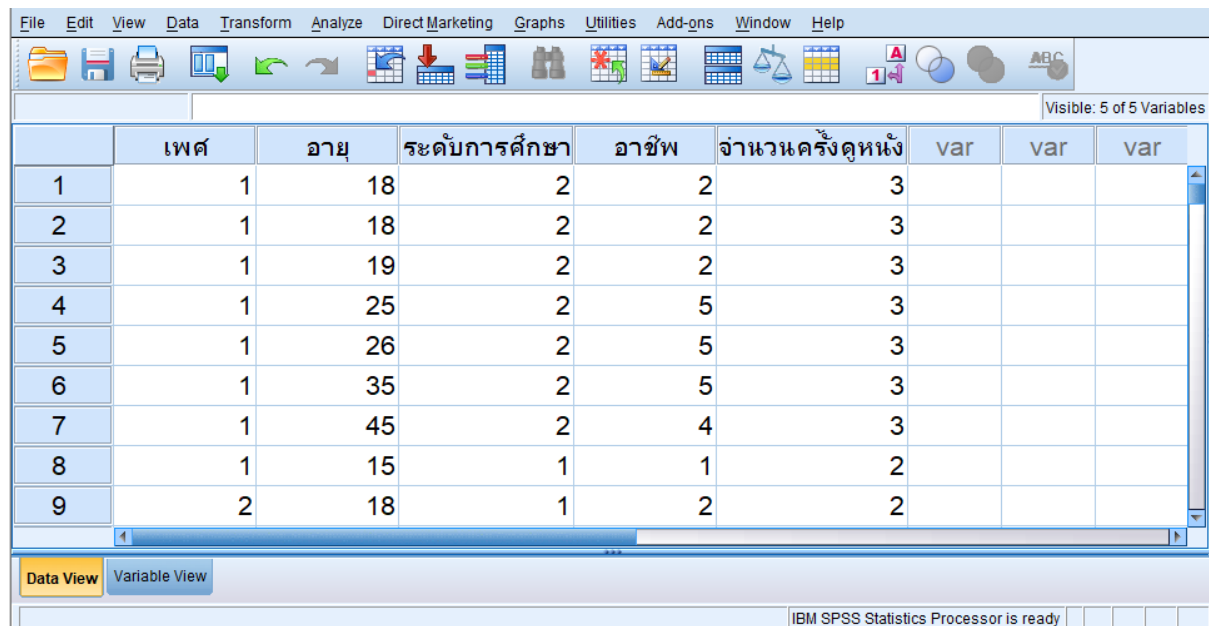
เมื่อกำหนดรายละเอียดทั้ง 5 ตัวแปรในหน้าต่าง Variable View และป้อนข้อมูลดังกล่าวของผู้ตอบแบบสอบถามในหน้าต่าง Data View เสร็จแล้วจะได้ดังภาพที่ 4.17



The screenshot shows the Variable View in IBM SPSS. The table below represents the data shown in the interface:

	Name	Type	Width	Decimals	Label	Values
1	เพศ	Numeric	8	0		{1, ชาย}...
2	อายุ	Numeric	8	0		None
3	ระดับการศึกษา	Numeric	8	0		{1, มัธยมศึกษา}...
4	อาชีพ	Numeric	8	0		{1, นักเรียน}...
5	จำนวนครั้งดูหนัง	Numeric	8	0	จำนวนครั้งดูหนังต่อเดือน	None
6						
7						
8						
9						
10						

At the bottom, the 'Variable View' tab is selected, and the status bar indicates 'IBM SPSS Statistics Processor is ready'.



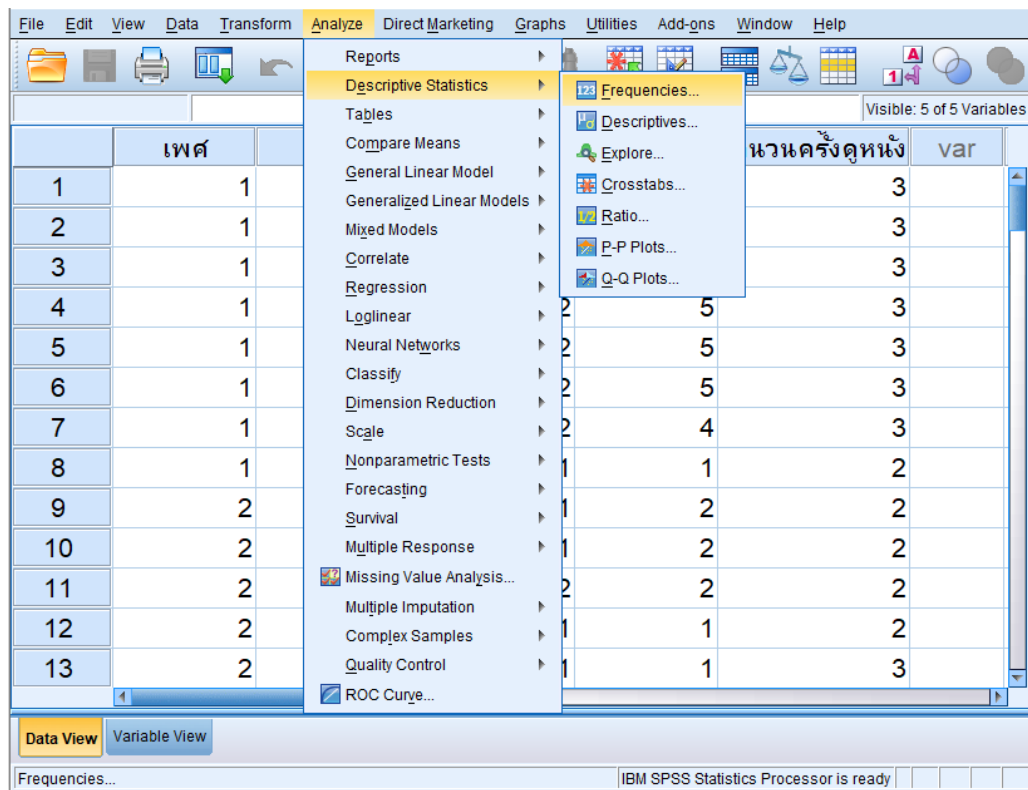
The screenshot shows the Data View in IBM SPSS. The table below represents the data shown in the interface:

	เพศ	อายุ	ระดับการศึกษา	อาชีพ	จำนวนครั้งดูหนัง	var	var	var
1	1	18	2	2	3			
2	1	18	2	2	3			
3	1	19	2	2	3			
4	1	25	2	5	3			
5	1	26	2	5	3			
6	1	35	2	5	3			
7	1	45	2	4	3			
8	1	15	1	1	2			
9	2	18	1	2	2			

At the bottom, the 'Data View' tab is selected, and the status bar indicates 'IBM SPSS Statistics Processor is ready'.

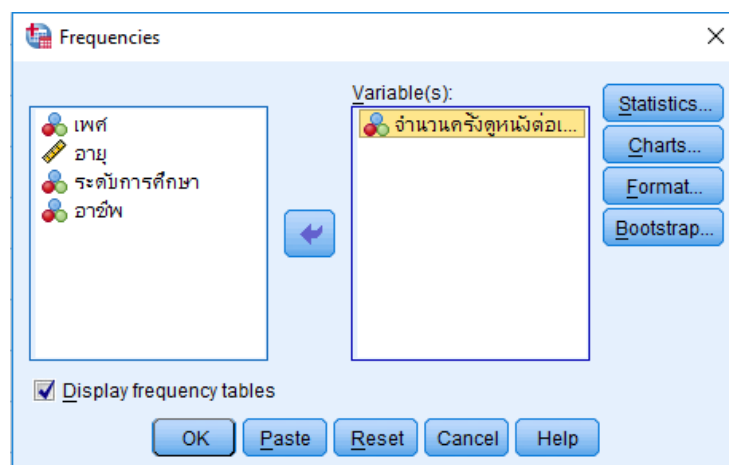
ภาพที่ 4.17 หน้าต่าง Variable View และ Data View เมื่อป้อนข้อมูลของผู้ตอบแบบสอบถาม

- (2) เมื่อป้อนข้อมูลทั้ง 5 ตัวแปรในโปรแกรม SPSS และตรวจสอบความถูกต้องแล้ว
ให้ใช้คำสั่ง **Analyze ➤ Descriptive Statistics ➤ Frequencies...** ดังภาพที่ 4.18 จะได้หน้าจอตั้ง
ภาพที่ 4.19



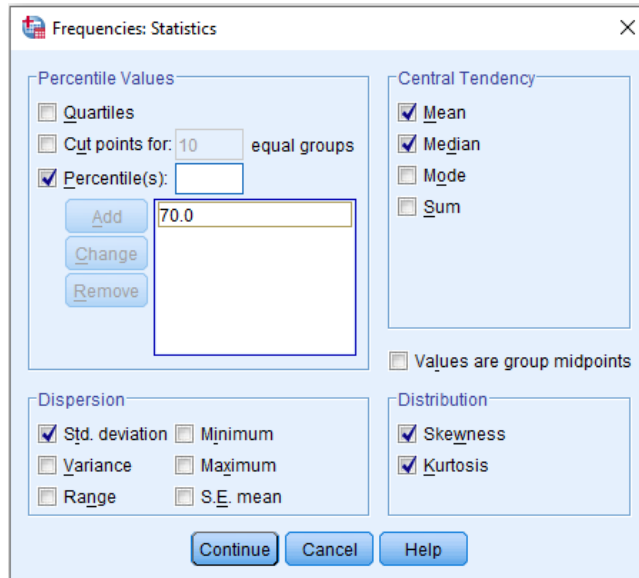
ภาพที่ 4.18 เมนูคำสั่ง Analyze ➤ Descriptive Statistics ➤ Frequencies...

- (3) เลือก ตัวแปรจำนวนครั้งดูหนังต่อเดือน (ตัวแปรเชิงปริมาณ) ที่อยู่ในกล่องด้านซ้ายมือไปใส่ใน
กล่อง Variable(s) ด้านขวามือ โดยทั่วไปผู้ใช้งานสามารถเลือกใส่ตัวแปรเชิงปริมาณได้ทีละหลายๆ ตัว



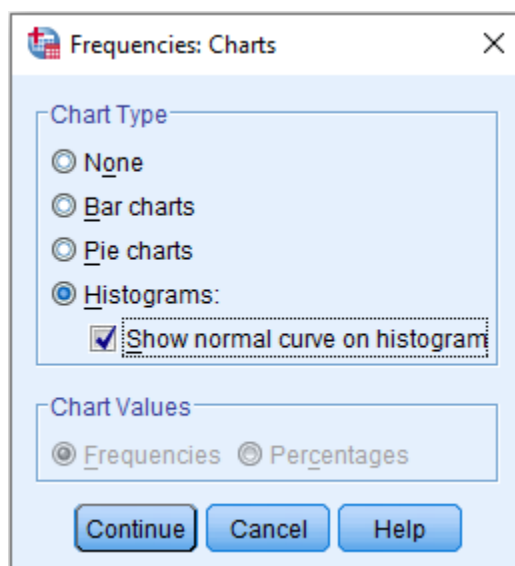
ภาพที่ 4.19 หน้าจอ Frequencies...

(4) กดคลิกเลือก **Statistics...** จะได้หน้าจอ ดังภาพที่ 4.20 ผู้ใช้สามารถเลือก Percentile Values, Central Tendency (สถิติที่วัดค่ากลางของข้อมูลเชิงปริมาณ), Dispersion (สถิติที่วัดการกระจายของข้อมูลเชิงปริมาณ) และ Distribution (สถิติที่แสดงการแจกแจงของข้อมูลเชิงปริมาณ เช่น Skewness แสดงค่าความเบ้ และ Kurtosis แสดงค่าความโด่ง) ได้ตามที่ต้องการ ในที่นี้ให้เลือกตามภาพที่ 4.20 แล้วกด Continue



ภาพที่ 4.20 หน้าจอ Frequencies : Statistics

(5) กดคลิกเลือก **Charts...** เพื่อกำหนดรูปแบบกราฟที่ต้องการนำเสนอ ในที่นี้ให้กดคลิกเลือก Histograms และ Show normal curve on histogram ดังภาพที่ 4.21 แล้วกด Continue และ OK จะได้ผลลัพธ์ดังภาพที่ 4.22



ภาพที่ 4.21 หน้าจอ Frequencies : Charts

Frequencies

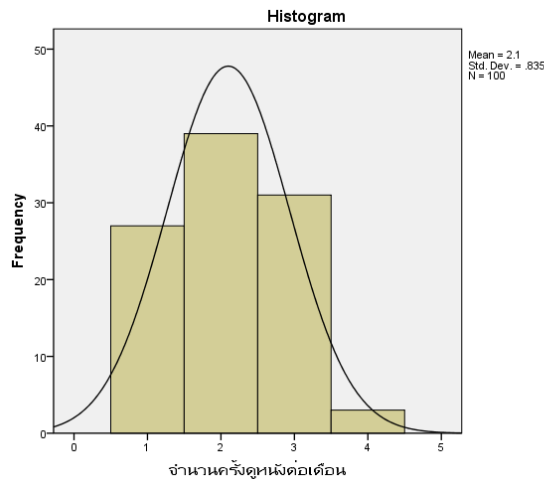
Statistics

จำนวนครั้งดูหนังต่อเดือน

N	Valid	100
	Missing	0
Mean		2.10
Median		2.00
Std. Deviation		.835
Skewness		.128
Std. Error of Skewness		.241
Kurtosis		-.907
Std. Error of Kurtosis		.478
Percentiles	70	3.00

จำนวนครั้งดูหนังต่อเดือน

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid 1	27	27.0	27.0	27.0
2	39	39.0	39.0	66.0
3	31	31.0	31.0	97.0
4	3	3.0	3.0	100.0
Total	100	100.0	100.0	



ภาพที่ 4.22 ผลลัพธ์การวิเคราะห์หาค่ากลางและค่าการกระจายด้วยคำสั่ง Frequencies

จากภาพที่ 4.22 ผู้ตอบแบบสอบถามไปดูหนังเฉลี่ย 2 ครั้งต่อเดือน (ค่า Mean เท่ากับ 2.10) เมื่อพิจารณาค่าเปอร์เซ็นต์ไทล์ที่ 70 มีค่าเท่ากับ 3.0 หมายถึง ผู้ตอบแบบสอบถามร้อยละ 70 ไปดูหนังไม่เกิน 3 ครั้งต่อเดือน และผู้ตอบแบบสอบถามอีกร้อยละ 30 ไปดูหนังเกิน 3 ครั้งต่อเดือน กราฟ Histogram แสดงให้เห็นว่า ตัวแปรจำนวนครั้งดูหนังต่อเดือน มีการแจกแจงแบบปกติหรือไม่

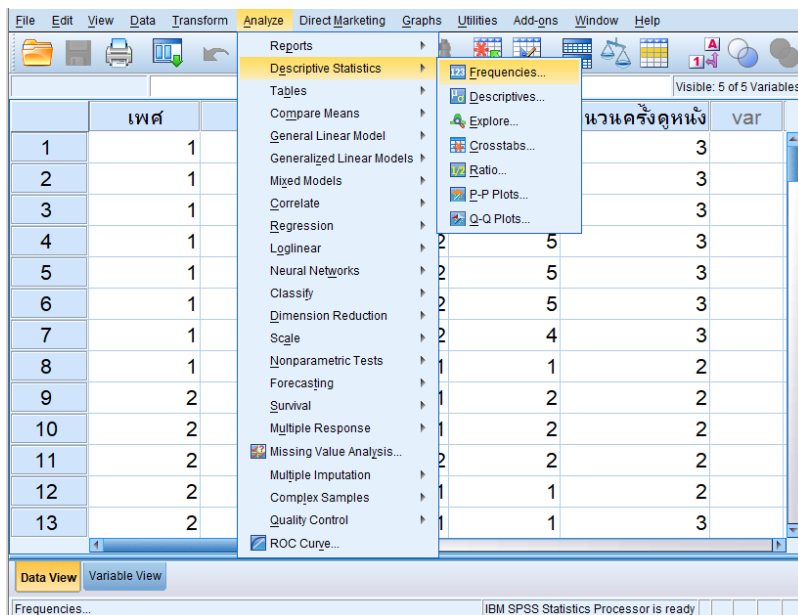
4.4.3.2 การวิเคราะห์หาค่ากลางและค่าการกระจายด้วยคำสั่ง Descriptives

(1) เมื่อป้อนข้อมูลในโปรแกรม SPSS แล้ว ให้ใช้คำสั่ง

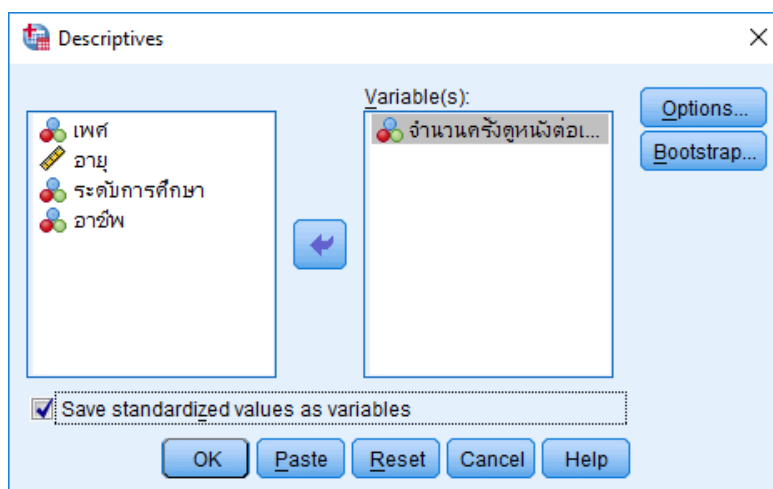
Analyze ➤ Descriptive Statistics ➤ Descriptives...

ดังภาพที่ 4.23 จะปรากฏหน้าต่างดังภาพที่

4.24 เลือก ตัวแปรจำนวนครั้งดูหนังต่อเดือน ที่อยู่ในกลุ่มด้านซ้ายมือไปใส่ในกลุ่ม Variable(s) ที่อยู่ด้านขวามือ และเลือก ☒ Save standardized values as variables เพื่อให้โปรแกรม SPSS คำนวณค่ามาตรฐานของตัวแปรจำนวนครั้งดูหนังต่อเดือน โดยโปรแกรมจะบันทึกค่ามาตรฐานนี้ใหม่โดยใช้ชื่อ Zจำนวนครั้งดูหนัง ต่อท้ายจากตัวแปรสุดท้ายในหน้าต่าง Data View ดังภาพที่ 4.25



ภาพที่ 4.23 เมนูคำสั่ง Analyze ➤ Descriptive Statistics ➤ Descriptives...



ที่ 4.24 หน้าจอ Descriptives

	เพศ	อายุ	ระดับการศึกษา	อาชีพ	จำนวนครั้งดูหนัง	Zจำนวนครั้งดูหนัง	var	var
1	1	18	2	2	3	1.08719		
2	1	18	2	2	3	1.08719		
3	1	19	2	2	3	1.08719		
4	1	25	2	5	3	1.08719		
5	1	26	2	5	3	1.08719		
6	1	35	2	5	3	1.08719		
7	1	45	2	4	3	1.08719		
8	1	15	1	1	2	-.10636		
9	2	18	1	2	2	-.10636		
10	2	17	1	2	2	-.10636		
11	2	18	2	2	2	-.10636		
12	2	15	1	1	2	-.10636		
13	2	16	1	1	3	1.08719		
14	2	15	1	1	3	1.08719		
15	2	14	1	1	3	1.08719		

ภาพที่ 4.25 ค่า Z จำนวนครั้งดูหนัง หรือค่ามาตรฐานของตัวแปรจำนวนครั้งดูหนังต่อเดือน

(2) กดเลือก **Options...** และเลือกคำสั่งตามหน้าจอตั้งภาพที่ 4.26 แล้วกด Continue และ OK จะได้ผลลัพธ์ดังภาพที่ 4.27

Descriptives: Options

☒ Mean ☐ Sum

Dispersion

☒ Std. deviation ☒ Minimum

☐ Variance ☒ Maximum

☐ Range ☐ S.E. mean

Distribution

☐ Kurtosis ☐ Skewness

Display Order

☒ Variable list

☐ Alphabetic

☐ Ascending means

☐ Descending means

Continue Cancel Help

ภาพที่ 4.26 หน้าจอ Descriptives : Options

Descriptives

Descriptive Statistics					
	N	Minimum	Maximum	Mean	Std. Deviation
จำนวนครั้งดูหนังต่อเดือน	100	1	4	2.10	.835
Valid N (listwise)	100				

ภาพที่ 4.27 ผลลัพธ์การวิเคราะห์หาค่ากลางและค่าการกระจายด้วยคำสั่ง Descriptives

4.5 การแจกแจงความถี่แบบ 2 ทางโดยใช้คำสั่ง Crosstabs

การแจกแจงความถี่มีความสำคัญต่อการนำเสนอข้อมูล เพื่อให้ นักวิจัยสามารถอธิบายข้อมูลได้อย่าง กระชับและถูกต้อง การแจกแจงความถี่สามารถนำเสนอในรูปแบบตาราง (Table) แผนภูมิวงกลม (Pie chart) แผนภูมิแท่ง (Bar chart) และกราฟฮิสโทแกรม (Histogram)

ในหัวข้อ 4.3 เป็นการสรุปลักษณะของข้อมูลโดยการแจกแจงความถี่ทางเดียว กรณีที่ต้องการแสดงความถี่หรือร้อยละของข้อมูลแยกตามตัวแปร 2 ตัว จะกำหนดให้ตัวแปรตัวหนึ่งอยู่ทางด้าน Row และอีกตัวแปรหนึ่งอยู่ทางด้าน Column ในแต่ละตัวแปรจะแบ่งเป็นหลายระดับ จะเรียกรายการนี้ว่า “**ตารางการแจกแจงความถี่ร่วม หรือ Crosstabs**” ตัวอย่างเช่น กรณีการสอบถามผู้บริโภคเรื่องพฤติกรรมในการดูหนัง ถ้า นักวิจัยต้องการสร้างตารางแจกแจงความถี่แยกตามเพศและอาชีพ สามารถเลือกให้ตัวแปรเพศอยู่ทางด้าน Row โดย Row จะแบ่งออกเป็น 2 ระดับ คือ เพศชาย และ เพศหญิง และเลือกให้ตัวแปรอาชีพอยู่ทางด้าน Column โดย Column จะแบ่งออกเป็น 5 ระดับ คือ นักเรียน, นิสิตนักศึกษา, รับราชการ, รัฐวิสาหกิจ และ พนักงานเอกชน

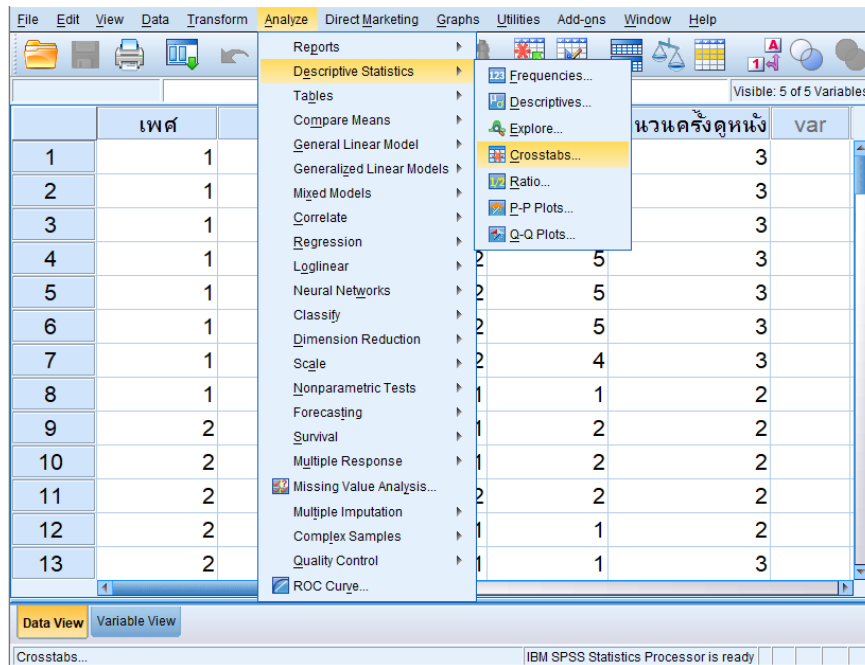
คำสั่งของ SPSS ในการสร้างตารางการแจกแจงความถี่ร่วม คือ

Analyze ➤ Descriptive Statistics ➤ Crosstabs...

ขั้นตอนการแจกแจงความถี่แบบ 2 ทางโดยใช้คำสั่ง Crosstabs

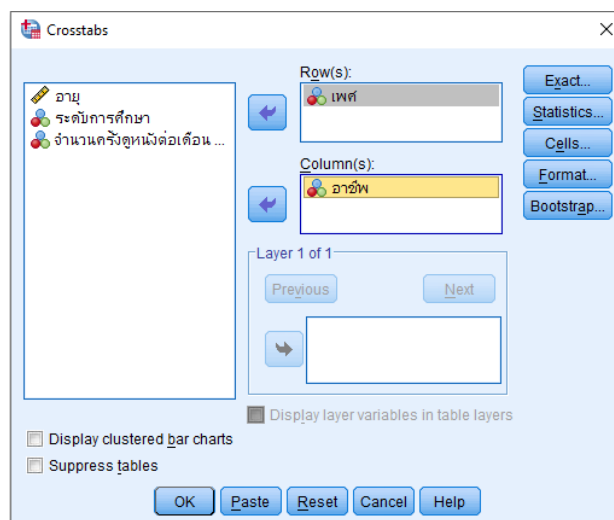
(1) เมื่อป้อนข้อมูลผู้ตอบแบบสอบถามในโปรแกรม SPSS เสร็จเรียบร้อยแล้ว ให้ใช้คำสั่ง

Analyze ➤ Descriptive Statistics ➤ Crosstabs... ดังภาพที่ 4.28



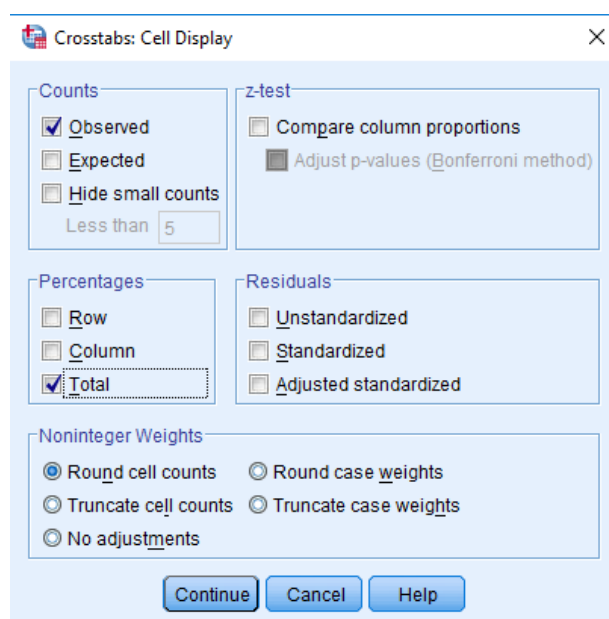
ภาพที่ 4.28 เมนูคำสั่ง Analyze ➤ Descriptive Statistics ➤ Crosstabs...

(2) จะได้หน้าจอ Crosstabs ดังภาพที่ 4.29 เลือกตัวแปรเพศ ใส่ในกล่อง Row(s) และตัวแปรอาชีพใส่ในกล่อง Column(s)



ภาพที่ 4.29 หน้าจอ Crosstabs...

(3) กดเลือก **Cells...** จะได้หน้าจอตั้งภาพที่ 4.30 ส่วนของ Counts เลือก Observed ส่วนของ Percentages เลือก Total ต่อจากนั้นกด Continue และ OK จะได้ผลลัพธ์ดังภาพที่ 4.31



ภาพที่ 4.30 หน้าจอ Crosstabs : Cell Display

Crosstabs

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
เพศ * อาชีพ	100	100.0%	0	.0%	100	100.0%

เพศ * อาชีพ Crosstabulation

			อาชีพ					Total
			นักเรียน	นิสิตนักศึกษา	รับราชการ	รัฐวิสาหกิจ	พนักงานบริษัทเอกชน	
เพศ ชาย	Count		10	15	4	10	14	53
	% of Total		10.0%	15.0%	4.0%	10.0%	14.0%	53.0%
หญิง	Count		15	13	9	5	5	47
	% of Total		15.0%	13.0%	9.0%	5.0%	5.0%	47.0%
Total	Count		25	28	13	15	19	100
	% of Total		25.0%	28.0%	13.0%	15.0%	19.0%	100.0%

ภาพที่ 4.31 ผลลัพธ์ตารางการแจกแจงความถี่ร่วมระหว่างตัวแปรเพศและอาชีพ

จากภาพที่ 4.31 เมื่อพิจารณาอาชีพในแต่ละคอลัมน์จะแยกเป็นเพศชายและเพศหญิง และจำนวนรวม เช่น ผู้ตอบแบบสอบถามทั้ง 100 คน เป็นนักเรียน 25 คน แบ่งเป็นเพศชาย 10 คน และเพศหญิง 15 คน

4.6 กรณีศึกษาการสร้างแบบสอบถาม การกำหนดตัวแปรและรหัสข้อมูลในโปรแกรม SPSS

4.6.1 ตัวอย่างการสร้างแบบสอบถามเพื่อทดสอบพฤติกรรมและความต้องการของผู้บริโภค

นายแฮร์รี พอตเตอร์ เป็นนักศึกษาปริญญาตรีที่ต้องการทำโครงการวิจัยเรื่อง การพัฒนาผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขมเพื่อผู้บริโภคที่ห่วงใยสุขภาพและผู้บริโภคทั่วไป จึงต้องการสอบถามความต้องการของผู้บริโภคเพื่อใช้เป็นข้อมูลในการตัดสินใจเลือกประเภทขนมขบเคี้ยวสำหรับพัฒนาผลิตภัณฑ์และกำหนดกลุ่มผู้บริโภคเป้าหมายที่สนใจในผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม นายแฮร์รี พอตเตอร์ จึงได้ทำการสร้างแบบสอบถามเพื่อใช้สำรวจผู้บริโภคจำนวน 100 คน ซึ่งประกอบด้วย ส่วนชี้แจงแบบสอบถาม และส่วนคำถามดังนี้

แบบสอบถาม

การสำรวจพฤติกรรมและความต้องการของผู้บริโภคต่อผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม

เรียน ผู้ตอบแบบสอบถาม

เรื่อง ขอความร่วมมือตอบแบบสอบถามเพื่อสำรวจพฤติกรรมและความต้องการของผู้บริโภคต่อผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม

คำชี้แจง แบบสอบถามนี้เป็นการสำรวจพฤติกรรมและความต้องการของผู้บริโภคต่อผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขมเพื่อใช้ประกอบการทำโครงการวิจัยของนายแฮร์รี พอตเตอร์ นักศึกษาปริญญาตรีของมหาวิทยาลัยฮอกวอตส์ จึงใคร่ขอความร่วมมือจากท่านกรุณาตอบแบบสอบถามให้สมบูรณ์

คำอธิบาย ผักโขมเป็นพืชที่มีโปรตีนสูง มีกรดอะมิโน วิตามิน และแร่ธาตุต่าง ๆ เช่น วิตามินเอ วิตามินบี ธาตุแคลเซียม ธาตุเหล็ก ธาตุแมกนีเซียม ธาตุโพแทสเซียม และแม้ผักโขมจะเป็นผักใบเขียวแต่มีปีตาแคโรทีนสูง โดยมีสารลูทีนและสารแซนโทน ซึ่งเป็นสารแคโรทีนอยด์อยู่เป็นจำนวนมาก นอกจากนั้นผักโขมยังมีเส้นใยอาหารมาก จึงช่วยการทำงานของระบบขับถ่าย และลดความเสี่ยงการเป็นมะเร็งกระเพาะอาหารได้

ขอขอบพระคุณยิ่ง
นายแฮร์รี พอตเตอร์
ผู้วิจัย

คำแนะนำ กรุณาทำเครื่องหมาย ✓ ลงในคำตอบที่ท่านเห็นว่าตรงกับความคิดของท่านมากที่สุด หรือเติมคำตอบให้สมบูรณ์

ส่วนที่ 1 ข้อมูลทั่วไปของผู้ตอบแบบสอบถาม

1. เพศ

() ชาย () หญิง

2. อายุ

() 15 - 20 ปี () 21 - 25 ปี () 26 - 30 ปี () 31 - 35 ปี

3. การศึกษา

() มัธยมศึกษา () ปริญญาตรี () ปริญญาโท () สูงกว่าปริญญาโท

4. อาชีพ

() นิสิต/นักศึกษา () แม่บ้าน () ข้าราชการ () รัฐวิสาหกิจ
() พนักงานเอกชน () ธุรกิจส่วนตัว () อื่นๆ โปรดระบุ.....

5. รายได้ต่อเดือน

() ต่ำกว่า 10,000 บาท () 10,000 - 15,000 บาท
() 15,001 - 20,000 บาท () สูงกว่า 20,000 บาท

ส่วนที่ 2 ข้อมูลเกี่ยวกับพฤติกรรมของผู้บริโภคต่อผลิตภัณฑ์ผักโขมและผลิตภัณฑ์ขนมขบเคี้ยว

1. ท่านรับประทานผลิตภัณฑ์ที่มีส่วนประกอบของผักโขมหรือไม่

() รับประทาน () ไม่รับประทาน (ข้ามไปตอบข้อ 3)

2. ผลิตภัณฑ์ที่มีส่วนประกอบของผักโขมใดที่ท่านรับประทาน (เลือกตอบได้มากกว่า 1 ข้อ)

() ผักโขมอบชีส () ข้าวผัดใส่ผักโขม () สเปกเก็ตตี้ผักโขม
() ขนมปังใส่ผักโขม () พายใส่ผักโขม () ซุปผักโขม
() คัสซาร์ดผักโขม () พิซซ่าหน้าผักโขม () สลัดผักโขม
() ไข่เจียวใส่ผักโขม () ต้มจืดผักโขม () อื่นๆ โปรดระบุ.....

3. ท่านรับประทานขนมขบเคี้ยวหรือไม่

() รับประทาน () ไม่รับประทาน

4. จำนวนเงินที่ท่านใช้ซื้อขนมขบเคี้ยวเพื่อรับประทานเอง.....บาท/ครั้ง

5. ท่านรับประทานขนมขบเคี้ยวต่อไปนี้ชนิดใดบ่อยที่สุด ให้เรียงลำดับ 3 อันดับแรก (โดยที่อันดับ 1 หมายถึง รับประทานบ่อยที่สุด)

- () ประเภทขนมปัง () ประเภทคุกกี้และบิสกิต
() ประเภทมันฝรั่งทอดกรอบ () ประเภทผักและผลไม้อบแห้ง
() ประเภทبوبพอง/อบกรอบ เช่น ข้าวพองอบกรอบ
() ประเภทถั่วและเมล็ดพืช

6. กิจกรรมที่ท่านมักทำร่วมกับการรับประทานขนมขบเคี้ยวมากที่สุด

- () ดูหนัง/ละคร () ระหว่างเดินทาง () ฟังเพลง
() อ่านหนังสือ () เล่นเกมส์ () อื่นๆ โปรดระบุ.....

7. ท่านให้ความสำคัญต่อปัจจัยในการเลือกซื้อขนมขบเคี้ยวอย่างไร

ปัจจัย	ระดับความสำคัญ				
	มากที่สุด	มาก	ปานกลาง	น้อย	น้อยที่สุด
คุณค่าทางโภชนาการ					
ความสวยงามของบรรจุภัณฑ์					
ราคา					
ปริมาณ/ขนาดบรรจุ					
รูปทรง/รูปร่างผลิตภัณฑ์					
ฟรีเซนเตอร์					
การส่งเสริมการขาย (ของแถม)					
อายุการเก็บรักษา					
เนื้อสัมผัส					
รสชาติ					
สีและลักษณะปรากฏ					
ใช้วัตถุดิบในประเทศ					
ขนาดชิ้น					
กลิ่น					
ไม่ใส่วัตถุกันเสีย					

ส่วนที่ 3 ข้อมูลเกี่ยวกับความต้องการของผู้บริโภคต่อผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม

1. หากมีการพัฒนาผลิตภัณฑ์ขนมขบเคี้ยวที่มีคุณค่าทางโภชนาการจากผักโขม ท่านสนใจหรือไม่

() สนใจ เพราะ.....

() ไม่สนใจ เพราะ..... (หยุดการตอบคำถาม)

2. ท่านต้องการให้ผลิตภัณฑ์ขนมขบเคี้ยวที่มีคุณค่าทางโภชนาการจากผักโขมเป็นผลิตภัณฑ์ขนมขบเคี้ยวประเภทใด ให้เรียงลำดับ 3 อันดับแรก (โดยที่อันดับ 1 หมายถึง ต้องการให้มากที่สุด)

() ประเภทขนมปัง

() ประเภทคุกกี้และบิสกิต

() ประเภทมันฝรั่งทอดกรอบ

() ประเภทผักและผลไม้อบแห้ง

() ประเภทอบพอง/อบกรอบ เช่น ข้าวพองอบกรอบ

() ประเภทถั่วและเมล็ดพืช

3. ท่านต้องการให้ผลิตภัณฑ์ขนมขบเคี้ยวที่มีคุณค่าทางโภชนาการจากผักโขม มีลักษณะปรากฏของผักโขมเป็นอย่างไร

() เป็นชิ้นของผักโขมขนาดเล็กกระจายตัวในผลิตภัณฑ์

() เป็นลักษณะผงผักโขมกระจายตัวในผลิตภัณฑ์

() เป็นครีมผักโขมเคลือบตัวผลิตภัณฑ์ด้านนอก

() อื่นๆ โปรดระบุ.....

4.6.2 ตัวอย่างการสร้างตัวแปรและกำหนดรหัสข้อมูลจากแบบสอบถามในโปรแกรม SPSS

จากตัวอย่างในข้อ 4.6.1 ซึ่งเป็นตัวอย่างการสร้างแบบสอบถามสำหรับสำรวจพฤติกรรมและความต้องการของผู้บริโภคต่อผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม เมื่อนักวิจัยต้องการป้อนข้อมูลในโปรแกรม SPSS เพื่อวิเคราะห์ข้อมูลทางสถิติ ไม่จำเป็นต้องกำหนดเลขที่ชุดของแบบสอบถามเนื่องจากโปรแกรมจะแสดงในรูปแบบเลขที่ Case ซึ่งเป็นเลขที่แสดงถึงแบบสอบถามได้อยู่แล้ว การสร้างตัวแปรและกำหนดรหัสสามารถเขียนได้ ดังนี้

ส่วนที่ 1 ข้อมูลทั่วไปของผู้ตอบแบบสอบถาม

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
1	A1	เพศ	0	1 = ชาย 2 = หญิง	เลือกได้ คำตอบเดียว
2	A2	อายุ	0	1 = 15-20 ปี 2 = 21-25 ปี 3 = 26-30 ปี 4 = 31-35 ปี	เลือกได้ คำตอบเดียว
3	A3	การศึกษา	0	1 = มัธยมศึกษา 2 = ปริญญาตรี 3 = ปริญญาโท 4 = สูงกว่าปริญญาโท	เลือกได้ คำตอบเดียว
4	A4	อาชีพ	0	1 = นิสิต/นักศึกษา 2 = แม่บ้าน 3 = ข้าราชการ 4 = รัฐวิสาหกิจ 5 = พนักงานเอกชน 6 = ธุรกิจส่วนตัว 7 = อื่นๆ	เลือกได้ คำตอบเดียว
5	A5	รายได้	0	1 = ต่ำกว่า 10,000 บาท 2 = 10,000-15,000 บาท 3 = 15,001-20,000 บาท 4 = สูงกว่า 20,000 บาท	เลือกได้ คำตอบเดียว

ส่วนที่ 2 ข้อมูลเกี่ยวกับพฤติกรรมของผู้บริโภคต่อผลิตภัณฑ์ผักโขมและผลิตภัณฑ์ขนมขบเคี้ยว

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
1	B1	ผักโขม	0	1 = รับประทาน 0 = ไม่รับประทาน	เลือกได้ คำตอบเดียว
2	B21	ผักโขมอบชีส	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	เลือกได้หลาย คำตอบโดย กำหนด จำนวนตัวแปร เท่ากับจำนวน ทางเลือก
	B22	ข้าวผัดใส่ผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B23	สปาเก็ตตี้ผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B24	ขนมปังใส่ผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B25	พายใส่ผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B26	ซูปรผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B27	คัสชั้วผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B28	พิซซ่าหน้าผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B29	สลัดผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B210	ไข่เจียวใส่ผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B211	ต้มจืดผักโขม	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
	B212	อื่นๆ	0	1 = ผู้ตอบเลือก 0 = ผู้ตอบไม่เลือก	
3	B3	ขนมขบเคี้ยว	0	1 = รับประทาน 0 = ไม่รับประทาน	เลือกได้ คำตอบเดียว
4	B4	จำนวนเงินที่ใช้ซื้อ ขนมขบเคี้ยวต่อครั้ง	2	ไม่ต้องกำหนด	ระบุจำนวน เงินที่แท้จริง

ส่วนที่ 2 (ต่อ)

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
5	B51	ชนิดขนมขบเคี้ยว อันดับ1	0	1 = ขนมปัง 2 = ลูกก๊วยชิต 3 = มันฝรั่งทอด 4 = ผักและผลไม้อบแห้ง 5 = อบปองอบกรอบ 6 = ถั่วและเมล็ดพืช	ให้ผู้ตอบ แบบสอบถาม จัดลำดับน้อย กว่าทางเลือก ที่ให้มา (ให้จัด 3 ลำดับจาก 6 ตัวเลือก) กำหนดให้มี จำนวนตัวแปร เท่ากับจำนวน ลำดับที่ให้จัด
	B52	ชนิดขนมขบเคี้ยว อันดับ2	0	1 = ขนมปัง 2 = ลูกก๊วยชิต 3 = มันฝรั่งทอด 4 = ผักและผลไม้อบแห้ง 5 = อบปองอบกรอบ 6 = ถั่วและเมล็ดพืช	
	B53	ชนิดขนมขบเคี้ยว อันดับ3	0	1 = ขนมปัง 2 = ลูกก๊วยชิต 3 = มันฝรั่งทอด 4 = ผักและผลไม้อบแห้ง 5 = อบปองอบกรอบ 6 = ถั่วและเมล็ดพืช	
6	B6	กิจกรรมที่ทำ ระหว่างรับประทานอาหาร ขนมขบเคี้ยว	0	1 = ดูหนัง/ละคร 2 = ระหว่างเดินทาง 3 = ฟังเพลง 4 = อ่านหนังสือ 5 = เล่นเกมส์ 6 = อื่นๆ	เลือกได้ คำตอบเดียว
7	B71	คุณค่าทาง โภชนาการ	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว

ส่วนที่ 2 (ต่อ)

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
7	B72	ความสวยงามของ บรรจุภัณฑ์	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B73	ราคา	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B74	ขนาดบรรจุ	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B75	รูปร่างผลิตภัณฑ์	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B76	ฟรีเซนเตอร์	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B77	การส่งเสริมการขาย	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว

ส่วนที่ 2 (ต่อ)

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
7	B78	อายุการเก็บรักษา	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B79	เนื้อสัมผัส	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B710	รสชาติ	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B711	สีและลักษณะ ปรากฏ	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B712	ใช้วัตถุดิบใน ประเทศ	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B713	ขนาดขึ้น	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว

ส่วนที่ 2 (ต่อ)

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
7	B714	กลืน	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว
	B715	ไม่ใส่วัตถุกันเสีย	0	5 = สำคัญมากที่สุด 4 = สำคัญมาก 3 = สำคัญปานกลาง 2 = สำคัญน้อย 1 = สำคัญน้อยที่สุด	เลือกได้ คำตอบเดียว

ส่วนที่ 3 ข้อมูลเกี่ยวกับความต้องการของผู้บริโภคต่อผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม

คำถามข้อที่	ชื่อตัวแปร (Name)	ความหมาย (Label)	จำนวน ทศนิยม (Decimals)	ค่ารหัสและความหมาย (Values)	ข้อสังเกต
1	C1	ผลิตภัณฑ์ขนมขบเคี้ยวจากผักโขม	0	1 = สนใจ 0 = ไม่สนใจ	เลือกได้ คำตอบเดียว
2	C21	ผลิตภัณฑ์ผักโขม อันดับ1	0	1 = ขนมปังผักโขม 2 = ลูกก๊อบบี้ผักโขม 3 = มันฝรั่งทอดผักโขม 4 = ผักโขมอบแห้ง 5 = ผักโขมอบพองกรอบ 6 = ถั่วและผักโขม	ให้ผู้ตอบ แบบสอบถาม จัดลำดับน้อย กว่าทางเลือก ที่ให้มา (ให้จัด 3 ลำดับจาก 6 ตัวเลือก) กำหนดให้มี จำนวนตัวแปร เท่ากับจำนวน ลำดับที่ให้จัด
	C22	ผลิตภัณฑ์ผักโขม อันดับ2	0	1 = ขนมปังผักโขม 2 = ลูกก๊อบบี้ผักโขม 3 = มันฝรั่งทอดผักโขม 4 = ผักโขมอบแห้ง 5 = ผักโขมอบพองกรอบ 6 = ถั่วและผักโขม	
	C23	ผลิตภัณฑ์ผักโขม อันดับ3	0	1 = ขนมปังผักโขม 2 = ลูกก๊อบบี้ผักโขม 3 = มันฝรั่งทอดผักโขม 4 = ผักโขมอบแห้ง 5 = ผักโขมอบพองกรอบ 6 = ถั่วและผักโขม	
3	C3	ลักษณะปรากฏของ ผักโขมในผลิตภัณฑ์	0	1 = เป็นชิ้น 2 = เป็นผง 3 = เป็นครีม 4 = อื่นๆ	เลือกได้ คำตอบเดียว



บทที่ 5

การทดสอบสมมติฐาน

สถิติอนุมาน เป็นวิชาสถิติที่ว่าด้วยกระบวนการสรุปข้อเท็จจริงเกี่ยวกับคุณลักษณะบางอย่างของประชากรโดยใช้ตัวอย่างที่สุ่มมาจากประชากร เรียกว่า ตัวอย่างสุ่ม (Random sample) ซึ่งใช้วิธีการสุ่ม (Random sampling) การอนุมานเชิงสถิติใช้ค่าที่คำนวณจากค่าสังเกตในตัวอย่างเป็นจำนวนมาก การแจกแจงความน่าจะเป็นของประชากรจากตัวอย่างที่สุ่มมา เรียกว่า การแจกแจงของตัวอย่างสุ่ม (Sampling distribution) ซึ่งมีหลายแบบ เช่น การแจกแจงแบบปกติ (Normal distribution), การแจกแจงแบบไคสแควร์ (Chi-square distribution), การแจกแจงแบบที (T distribution) และ การแจกแจงแบบเอฟ (F distribution) ทั้งนี้การแจกแจงของตัวอย่างสุ่มที่สำคัญที่สุดชนิดหนึ่ง คือ การแจกแจงแบบปกติ เนื่องจากวิธีการวิเคราะห์ข้อมูลทางสถิติส่วนใหญ่ที่นักวิจัยนิยมใช้กันนั้นมีข้อสมมติฐานเบื้องต้นว่า ข้อมูลของตัวอย่างที่สุ่มมา มีการแจกแจงแบบปกติ ลักษณะของกราฟการแจกแจงเป็นรูปประฆังที่มีสมมาตร มีความโค้งพอดี

ในการทำวิจัยไม่ว่าจะใช้ข้อมูลเชิงคุณภาพหรือเชิงปริมาณ หรือข้อมูลทั้ง 2 ประเภทเพื่อหาคำตอบในสิ่งที่นักวิจัยสนใจ เช่น อาชีพมีผลต่อการยอมรับผลิตภัณฑ์หรือไม่ การใช้อุณหภูมิอบที่แตกต่างกันมีผลต่อสีของกล้วยตากหรือไม่ จำเป็นต้องมีการทดสอบสมมติฐานเพื่อสรุปคำตอบดังกล่าว เพื่อหาสาเหตุ หรือเพื่อการทดสอบความสัมพันธ์ เงื่อนไขของการทดสอบสมมติฐาน คือ การทดสอบสมมติฐานทางสถิติโดยใช้ข้อมูลตัวอย่าง ประชากร หรือตัวแปรที่ต้องการทดสอบจะต้องมีการแจกแจงแบบปกติหรือใกล้เคียงแบบปกติ

5.1 ความหมายของสมมติฐาน (Hypothesis)

สมมติฐานวิจัย คือ ข้อสมมติ ข้อความเฉพาะที่ผู้วิจัยคาดคะเนคำตอบไว้ล่วงหน้าอย่างมีเหตุผลหรือข้อค้นพบในเรื่องที่ทำการศึกษาวิจัย โดยคำตอบที่คาดคิดไว้ล่วงหน้านี้อาจเกิดจากทฤษฎีที่เกี่ยวข้องหรือจากประสบการณ์ของผู้วิจัยเอง เพื่อใช้เป็นแนวทางในการค้นคว้าวิจัย รวบรวมข้อมูลและวิเคราะห์ข้อมูลว่าสิ่งที่ผู้วิจัยคาดไว้นั้นเป็นไปตามที่คาดไว้หรือไม่ เช่น จากประสบการณ์ของอาจารย์ผู้สอนคาดว่า การเข้าห้องเรียนมีผลต่อความพึงพอใจของนิสิตในการเรียน หรือ นักพัฒนาผลิตภัณฑ์คาดว่า การเปลี่ยนวัตถุดิบไม่มีผลต่อคะแนนความชอบของผู้บริโภค

สมมติฐานมีบทบาทสำคัญมากในการดำเนินงานวิจัย เนื่องจากเป็นตัวชี้้นำในการค้นหาข้อเท็จจริง การดำเนินงานวิจัยจึงเริ่มต้นด้วยการตั้งสมมติฐานงานวิจัย ประโยชน์ของสมมติฐานงานวิจัย คือ ช่วยจำกัดขอบเขตของการวิจัย ช่วยทำให้เห็นความสัมพันธ์ของข้อมูลที่น่ามาทดสอบสมมติฐาน ช่วยชี้แบบแผนในการวิจัยให้ผู้วิจัยเลือกเครื่องมือในการเก็บรวบรวมและวิเคราะห์ข้อมูล สมมติฐานช่วยให้ผู้วิจัยคาดคะเนและทดสอบเฉพาะตัวแปรที่เกี่ยวข้องกับปัญหาที่สนใจ ไม่จำเป็นต้องศึกษาและทดสอบตัวแปรอื่นที่ไม่เกี่ยวข้องจึงทำให้การวิจัยประหยัดเวลาและค่าใช้จ่าย

5.2 ประเภทของสมมติฐาน

สมมติฐานโดยทั่วไปแบ่งออกเป็น 2 ประเภท คือ

5.2.1 สมมติฐานการวิจัย (Research hypothesis) เป็นข้อความที่แสดงความสัมพันธ์ของตัวแปรที่ศึกษาซึ่งเป็นคำตอบของนักวิจัยที่ได้คาดหมายไว้ และได้จากการศึกษาเอกสารและงานวิจัยที่เกี่ยวข้อง โดยการตั้งสมมติฐานจะขึ้นอยู่กับทฤษฎี ผลการวิจัย และแนวความคิดของผู้วิจัยที่ได้ศึกษาค้นคว้าเบื้องต้น โดยจะเขียนในลักษณะที่เข้าใจง่าย สามารถสื่อความหมายได้โดยตรง

5.2.2 สมมติฐานทางสถิติ (Statistical hypothesis) เป็นข้อความเกี่ยวกับค่าพารามิเตอร์ที่นักวิจัยกำหนดขึ้นตามหลักการของสถิติที่ต้องการทดสอบ เป็นข้อความที่ตั้งขึ้นเพื่อทดสอบสมมติฐานที่ผู้วิจัยตั้งไว้เป็นจริงหรือไม่ สมมติฐานทางสถิติประกอบด้วย 2 ส่วน คือ

(1) **สมมติฐานว่าง หรือ สมมติฐานหลัก (Null hypothesis)** ใช้สัญลักษณ์ H_0 เป็นสมมติฐานที่ระบุถึงการไม่มีความสัมพันธ์กัน หรือ ไม่มีความแตกต่างของค่าพารามิเตอร์ เช่น

- นิสิตที่เรียนออนไลน์ได้คะแนนสอบไม่แตกต่างกับนิสิตที่เรียนในห้องเรียน
- ผู้ทดสอบให้คะแนนความชอบเค้กที่ใช้สารให้ความหวานไม่แตกต่างกับเค้กที่ใช้น้ำตาล
- อุณหภูมิในการอบที่แตกต่างกันไม่ส่งผลต่อค่าความสว่างของขนมปัง

(2) **สมมติฐานทางเลือก หรือ สมมติฐานรอง (Alternative hypothesis)** ใช้สัญลักษณ์ H_1 หรือ H_a เป็นสมมติฐานที่ระบุค่าพารามิเตอร์ตามที่ผู้วิจัยคาดหมายไว้หรือตามสมมติฐานการวิจัย เป็นสมมติฐานที่มีลักษณะแตกต่างจากสมมติฐานหลักที่จะทดสอบ กล่าวได้ว่า สมมติฐานรอง คือ สมมติฐานหลักที่ถูกปฏิเสธ โดยสมมติฐานรองจะบอกถึงความแตกต่างในลักษณะ มากกว่า หรือ น้อยกว่า หรือ แตกต่างกัน ขึ้นอยู่กับการตั้งสมมติฐาน สมมติฐานรองแบ่งออกเป็น 2 ประเภท คือ

- **สมมติฐานรองแบบไม่มีทิศทาง (Non-directional alternative hypothesis)** หรือ การทดสอบสมมติฐานแบบสองหาง (Two-tailed test) เป็นสมมติฐานที่นักวิจัยสนใจศึกษาว่าค่าพารามิเตอร์มีความแตกต่างหรือไม่ โดยไม่สนใจว่าแตกต่างไปในทิศทางมากกว่าหรือน้อยกว่า เช่น

- นิสิตที่เรียนออนไลน์ได้คะแนนสอบแตกต่างกับนิสิตที่เรียนในห้องเรียน
- ผู้ทดสอบให้คะแนนความชอบเค้กที่ใช้สารให้ความหวานแตกต่างกับเค้กที่ใช้น้ำตาล
- อุณหภูมิในการอบที่แตกต่างกันส่งผลต่อค่าความสว่างของขนมปัง

- สมมติฐานรองแบบมีทิศทาง (Directional alternative hypothesis) หรือ การทดสอบสมมติฐานแบบหางเดียว (One-tailed test) เป็นสมมติฐานที่นักวิจัยสนใจศึกษาว่าค่าพารามิเตอร์มีความแตกต่างไปในทิศทางมากกว่าหรือน้อยกว่า เช่น

- นิสิตที่เรียนออนไลน์ได้คะแนนสอบน้อยกว่านิสิตที่เรียนในห้องเรียน
- ผู้ทดสอบให้คะแนนความชอบเค้กที่ใช้สารให้ความหวานน้อยกว่าเค้กที่ใช้น้ำตาล
- มันทิ้งที่ทอดแบบน้ำมันท่วมมีปริมาณไขมันมากกว่าทอดแบบวิธีสุญญากาศ

5.3 ประเภทของการทดสอบสมมติฐาน

การทดสอบสมมติฐาน (Hypothesis testing) คือ การทดสอบเพื่อที่จะสรุปว่าสิ่งที่ผู้วิจัยคาดไว้หรือเชื่อเป็นจริงหรือไม่ จึงต้องมีการรวบรวมข้อมูลจริงเพื่อมาทดสอบ บางครั้งเรียกการทดสอบสมมติฐานว่า การทดสอบนัยสำคัญทางสถิติ (Significant testing)

ในการทดสอบสมมติฐานทางสถิติ จะต้องมียสมมติฐานหลัก (H_0) และสมมติฐานรอง (H_1) ควบคู่กันไป และข้อความของสมมติฐานหลักและสมมติฐานรองจะต้องแยกกันโดยเด็ดขาดเพื่อประโยชน์ในการสรุปผล ในการทดสอบสมมติฐานทางสถิติจะมีผลการตัดสินใจ 2 ลักษณะ คือ หากนักวิจัยมีหลักฐานหรือข้อพิสูจน์เพียงพอที่จะยอมรับหรือทำให้เชื่อว่าสมมติฐานหลักถูกต้อง นักวิจัยจะยอมรับสมมติฐานหลัก (Accept H_0) ในทางตรงข้ามถ้าข้อพิสูจน์นั้นแสดงให้เห็นว่าสมมติฐานหลักไม่ถูกต้อง นักวิจัยจะปฏิเสธสมมติฐานหลัก (Reject H_0) และจะยอมรับสมมติฐานรอง (Accept H_1) กล่าวสรุปได้ ดังนี้

(1) การยอมรับสมมติฐานหลัก (Accept H_0) เป็นผลของการยอมรับสมมติฐานหลัก (H_0) ที่ตั้งไว้ โดยถือว่าความแตกต่างระหว่างค่าสถิติจากกลุ่มตัวอย่างกับค่าพารามิเตอร์ที่คาดหวังไว้ในสมมติฐานหลักมีเพียงเล็กน้อย เป็นความแตกต่างโดยบังเอิญไม่มีนัยสำคัญ อาจเกิดจากความคลาดเคลื่อนในการสุ่มตัวอย่าง หรือความคลาดเคลื่อนจากการเก็บรวบรวมข้อมูล ผลการทดสอบนี้เรียกว่า “การทดสอบที่ไม่มีนัยสำคัญ (Non-significant)”

(2) การปฏิเสธสมมติฐานหลัก (Reject H_0) เป็นผลของการปฏิเสธสมมติฐานหลัก (H_0) ที่ตั้งไว้ โดยความแตกต่างระหว่างค่าสถิติจากกลุ่มตัวอย่างกับค่าพารามิเตอร์ที่คาดหวังไว้ในสมมติฐานหลักมีมากจนถือเป็นความแตกต่างที่แท้จริงหรือแตกต่างอย่างมีนัยสำคัญ ผลการทดสอบนี้เรียกว่า “การทดสอบที่มีนัยสำคัญ (Significant)”

ประเภทของการทดสอบสมมติฐาน แบ่งเป็น 2 ประเภทใหญ่ๆ คือ

(1) การทดสอบสมมติฐานแบบสองหาง (Two-tailed test หรือ Two-side test)

การทดสอบสมมติฐานแบบสองหาง บางครั้งเรียกว่า การทดสอบสมมติฐานแบบสองหาง ในการเขียนหรือตั้งสมมติฐานหลัก (H_0) จะใช้เครื่องหมายเท่ากับ (=) เสมอ และสมมติฐานรอง (H_1) จะใช้เครื่องหมายไม่เท่ากับ ($\neq$) เสมอ

ตัวอย่างที่ 1 นักวิจัยต้องการทราบว่า การใช้ปริมาณน้ำตาลที่แตกต่างกัน (30, 40 และ 50%) มีผลต่อค่าสี L^* ของทุเรียนกวนหรือไม่ นักวิจัยสามารถเขียนสมมติฐาน H_0 และ H_1 ได้หลายแบบ ดังนี้

H_0 : ปริมาณน้ำตาลไม่มีอิทธิพลต่อค่าเฉลี่ยสี L^* ของทุเรียนกวน

H_1 : ปริมาณน้ำตาลมีอิทธิพลต่อค่าเฉลี่ยสี L^* ของทุเรียนกวน

H_0 : ปริมาณน้ำตาลไม่มีผลต่อค่าเฉลี่ยสี L^* ของทุเรียนกวน

H_1 : ปริมาณน้ำตาลมีผลต่อค่าเฉลี่ยสี L^* ของทุเรียนกวน

H_0 : ค่าเฉลี่ยสี L^* ของทุเรียนกวนทุกสิ่งทดลองไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยสี L^* ของทุเรียนกวนอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

หรือเขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

H_0 : $\mu_{30} = \mu_{40} = \mu_{50}$

H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่ ; $i \neq j$

ตัวอย่างที่ 2 นักวิจัยต้องการทราบว่า การใช้อุณหภูมิในการอบที่แตกต่างกัน (50, 60 และ 70 °C) ในเวลาอบที่เท่ากัน มีผลต่อปริมาณความชื้นของกล้วยตากหรือไม่ นักวิจัยสามารถเขียนสมมติฐาน H_0 และ H_1 ได้หลายแบบ ดังนี้

H_0 : อุณหภูมิในการอบไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณความชื้นของกล้วยตาก

H_1 : อุณหภูมิในการอบมีอิทธิพลต่อค่าเฉลี่ยปริมาณความชื้นของกล้วยตาก

H_0 : อุณหภูมิในการอบไม่มีผลต่อค่าเฉลี่ยปริมาณความชื้นของกล้วยตาก

H_1 : อุณหภูมิในการอบมีผลต่อค่าเฉลี่ยปริมาณความชื้นของกล้วยตาก

- H_0 : ค่าเฉลี่ยปริมาณความชื้นของกล้วยตากทุกสิ่งทดลองไม่แตกต่างกัน
 H_1 : ค่าเฉลี่ยปริมาณความชื้นของกล้วยตากอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

หรือเขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

- H_0 : $\mu_{50} = \mu_{60} = \mu_{70}$
 H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่ ; $i \neq j$

ตัวอย่างที่ 3 นักวิจัยต้องการทราบว่า ผู้ทดสอบ 30 คนให้คะแนนความชอบน้ำส้มสองยี่ห้อ (ยี่ห้อ A และ B) แตกต่างกันหรือไม่ นักวิจัยสามารถเขียนสมมติฐาน H_0 และ H_1 ได้หลายแบบ ดังนี้

- H_0 : ยี่ห้อน้ำส้มไม่มีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบ
 H_1 : ยี่ห้อน้ำส้มมีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบ

- H_0 : ยี่ห้อน้ำส้มไม่มีผลต่อค่าเฉลี่ยคะแนนความชอบ
 H_1 : ยี่ห้อน้ำส้มมีผลต่อค่าเฉลี่ยคะแนนความชอบ

- H_0 : ค่าเฉลี่ยคะแนนความชอบน้ำส้มยี่ห้อ A ไม่แตกต่างกับยี่ห้อ B
 H_1 : ค่าเฉลี่ยคะแนนความชอบน้ำส้มยี่ห้อ A แตกต่างกับยี่ห้อ B

หรือเขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

- H_0 : $\mu_A = \mu_B$
 H_1 : $\mu_A \neq \mu_B$

ตัวอย่างที่ 4 นักวิจัยต้องการทราบว่า สูตรเค้กที่พัฒนาขึ้นมาใหม่มีปริมาณไขมันเท่ากับ 20% หรือไม่ นักวิจัยสามารถเขียนสมมติฐาน H_0 และ H_1 ได้หลายแบบ ดังนี้

- H_0 : สูตรเค้กที่พัฒนาขึ้นมาใหม่มีค่าเฉลี่ยปริมาณไขมันเท่ากับ 20%
 H_1 : สูตรเค้กที่พัฒนาขึ้นมาใหม่มีค่าเฉลี่ยปริมาณไขมันไม่เท่ากับ 20%

หรือเขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

- H_0 : $\mu = 20\%$
 H_1 : $\mu \neq 20\%$

(2) การทดสอบสมมติฐานแบบหางเดียว (One-tailed test หรือ One-side test)

การทดสอบสมมติฐานแบบหางเดียว บางครั้งเรียกว่า การทดสอบสมมติฐานแบบหางเดียว ในการเขียนหรือตั้งสมมติฐานหลัก (H_0) และสมมติฐานรอง (H_1) มี 2 กรณี ดังแสดงในตารางที่ 5.1

ตารางที่ 5.1 การกำหนดเครื่องหมายการทดสอบสมมติฐานแบบหางเดียว

สิ่งที่นักวิจัยคาดการณ์ไว้	สมมติฐานหลัก (H_0)	สมมติฐานรอง (H_1)
กรณีที่ 1 นักวิจัยคาดว่าสิ่งที่เกิดขึ้นมีทิศทางมากกว่า	$\leq$	$>$
กรณีที่ 2 นักวิจัยคาดว่าสิ่งที่เกิดขึ้นมีทิศทางน้อยกว่า	$\geq$	$<$

กรณีที่ 1 หากนักวิจัยคาดว่าสิ่งที่เกิดขึ้น มีทิศทางมากกว่า ให้กำหนดเครื่องหมายของสมมติฐานรอง (H_1) เป็น เครื่องหมายมากกว่า

ตัวอย่างที่ 1 อาจารย์คาดว่า ค่าเฉลี่ยคะแนนความพึงพอใจของนิสิตต่อการสอนมากกว่า 4 คะแนน (1 = พอใจน้อยที่สุด และ 5 = พอใจมากที่สุด) สิ่งทีอาจารย์คาดไว้มี ทิศทางมากกว่า ดังนั้นจึงเขียนสมมติฐานได้ ดังนี้

H_0 : ค่าเฉลี่ยคะแนนความพึงพอใจของนิสิตต่อการสอน ≤ 4 คะแนน

H_1 : ค่าเฉลี่ยคะแนนความพึงพอใจของนิสิตต่อการสอน > 4 คะแนน

ตัวอย่างที่ 2 นักวิจัยคาดว่า มันฝรั่งที่ทอดแบบน้ำมันท่วมมีปริมาณไขมันมากกว่าทอดแบบวิธีสุญญากาศ สิ่งทีนักวิจัยคาดไว้มี ทิศทางมากกว่า ดังนั้นจึงเขียนสมมติฐานได้ ดังนี้

H_0 : มันฝรั่งที่ทอดแบบน้ำมันท่วมมีค่าเฉลี่ยปริมาณไขมัน $\leq$ มันฝรั่งที่ทอดแบบวิธีสุญญากาศ

H_1 : มันฝรั่งที่ทอดแบบน้ำมันท่วมมีค่าเฉลี่ยปริมาณไขมัน $>$ มันฝรั่งที่ทอดแบบวิธีสุญญากาศ

กรณีที่ 2 หากนักวิจัยคาดว่าสิ่งที่เกิดขึ้น มีทิศทางน้อยกว่า ให้กำหนดเครื่องหมายของสมมติฐานรอง (H_1) เป็น เครื่องหมายน้อยกว่า

ตัวอย่างที่ 1 นักวิจัยคาดว่า ผู้บริโภควัยทำงานมีความถี่ในการไปดูหนังต่อสัปดาห์น้อยกว่าผู้บริโภควัยเรียน สิ่งทีนักวิจัยคาดไว้มี ทิศทางน้อยกว่า ดังนั้นจึงเขียนสมมติฐานได้ ดังนี้

H_0 : ผู้บริโภควัยทำงานมีความถี่ในการไปดูหนังต่อสัปดาห์ $\geq$ ผู้บริโภควัยเรียน

H_1 : ผู้บริโภควัยทำงานมีความถี่ในการไปดูหนังต่อสัปดาห์ $<$ ผู้บริโภควัยเรียน

ตัวอย่างที่ 2 นักวิจัยคาดว่า เค้กที่ใช้สารให้ความหวานแทนที่น้ำตาลมีปริมาณแคลอรีน้อยกว่า 250 kcal สิ่งที่นักวิจัยคาดไว้มี ทิศทางน้อยกว่า ดังนั้นจึงเขียนสมมติฐานได้ ดังนี้

H_0 : เค้กที่ใช้สารให้ความหวานแทนที่น้ำตาลมีปริมาณแคลอรี ≥ 250 kcal

H_1 : เค้กที่ใช้สารให้ความหวานแทนที่น้ำตาลมีปริมาณแคลอรี < 250 kcal

สรุปหลักเกณฑ์ในการเขียนสมมติฐานหลักและสมมติฐานรอง

นักวิจัยต้องพิจารณาถึงสิ่งที่คาดการณ์ไว้ โดยให้เขียนสมมติฐานรอง (H_1) ก่อน จากนั้นเขียนสมมติฐานหลัก (H_0) ในทิศทางตรงข้าม แบ่งการพิจารณาออกได้เป็น 3 กรณี

กรณีที่ 1 สิ่งที่นักวิจัยคาดการณ์ไว้หรือสนใจ คือ ความแตกต่างแบบไม่สนใจทิศทาง ให้เขียนสมมติฐานรอง (H_1) โดยใช้เครื่องหมายไม่เท่ากับ ($\neq$) และเขียนสมมติฐานหลัก (H_0) โดยใช้เครื่องหมายเท่ากับ ($=$)

กรณีที่ 2 สิ่งที่นักวิจัยคาดการณ์ไว้หรือสนใจ คือ ความแตกต่างแบบมีทิศทางในด้านบวก หรือมากกว่า ให้เขียนสมมติฐานรอง (H_1) โดยใช้เครื่องหมายมากกว่า ($>$) และเขียนสมมติฐานหลัก (H_0) โดยใช้เครื่องหมายน้อยกว่าหรือเท่ากับ ($\leq$)

กรณีที่ 3 สิ่งที่นักวิจัยคาดการณ์ไว้หรือสนใจ คือ ความแตกต่างแบบมีทิศทางในด้านลบ หรือน้อยกว่า ให้เขียนสมมติฐานรอง (H_1) โดยใช้เครื่องหมายน้อยกว่า ($<$) และเขียนสมมติฐานหลัก (H_0) โดยใช้เครื่องหมายมากกว่าหรือเท่ากับ ($\geq$)

5.4 ความคลาดเคลื่อนในการทดสอบสมมติฐานทางสถิติและระดับนัยสำคัญ

ในการทดสอบสมมติฐานทางสถิติเพื่อหาความสัมพันธ์หรือสาเหตุ หรือเพื่อทดสอบว่าสิ่งที่ผู้วิจัยคาดไว้จะเป็นจริงหรือไม่นั้น อาจเกิดความคลาดเคลื่อนหรือความผิดพลาดในการสรุปผลขึ้นได้ โดยเฉพาะอย่างยิ่งกรณีถ้าข้อมูลตัวอย่างที่นำมาใช้ในการทดสอบไม่มีคุณภาพเพียงพอ ความคลาดเคลื่อนในการทดสอบสมมติฐานทางสถิติ แบ่งออกเป็น 2 ประเภท ดังนี้

5.4.1 ความคลาดเคลื่อนประเภทที่ 1 (Type I error) คือ ความผิดพลาดที่เกิดขึ้นเนื่องจากผู้วิจัยสรุปว่า สมมติฐานหลัก (H_0) ไม่จริง ทั้งที่ในความเป็นจริงนั้นสมมติฐาน H_0 จริง ค่าความน่าจะเป็นที่จะเกิดความผิดพลาดหรือโอกาสที่ผู้วิจัยจะสรุปผิดกำหนดด้วย ค่า α (Alpha) หรือ ค่าระดับนัยสำคัญ (Level of significance) โดยปกติในงานวิจัยและพัฒนาผลิตภัณฑ์จะกำหนดให้ α มีค่าเท่ากับ 0.05 และ 0.01

ตัวอย่าง นายณเดชน์คาดว่าจะมีนักเรียนขี่จักรยานผ่านหน้าบ้านจำนวน 4 คน จึงตั้งสมมติฐานงานวิจัย ดังนี้

H_0 : มีนักเรียนขี่จักรยานผ่านหน้าบ้าน = 4 คน

H_1 : มีนักเรียนขี่จักรยานผ่านหน้าบ้าน $\neq$ 4 คน

เมื่อทำการสรุปผลการทดสอบสมมติฐาน :



นายณเดชน์ สรุป ปฏิเสธ H_0 นั่นคือ มีนักเรียนขี่จักรยานผ่านหน้าบ้านไม่เท่ากับ 4 คน ทั้งที่ความเป็นจริงมีนักเรียนขี่จักรยานผ่านหน้าบ้าน 4 คน กรณีนี้เป็นความคลาดเคลื่อนประเภทที่ 1 เนื่องจากผู้วิจัยสรุปว่า สมมติฐานหลัก (H_0) ไม่จริง ทั้งที่ในความเป็นจริงนั้นสมมติฐาน H_0 จริง

5.4.2 ความคลาดเคลื่อนประเภทที่ 2 (Type II error) คือ ความผิดพลาดที่เกิดขึ้นเนื่องจากผู้วิจัยยอมรับสมมติฐานหลัก (H_0) ทั้งที่ในความเป็นจริงนั้นสมมติฐาน H_0 ไม่จริง ค่าความน่าจะเป็นที่จะเกิดความผิดพลาดหรือโอกาสที่ผู้วิจัยจะสรุปผิดกำหนดด้วย ค่า β (Beta)

ตัวอย่าง นางสาวญาญ่าคาดว่าจะมีแอปเปิ้ลในตะกร้ามากกว่า 5 ลูก จึงตั้งสมมติฐานงานวิจัย ดังนี้

H_0 : มีแอปเปิ้ลในตะกร้า $\leq$ 5 ลูก

H_1 : มีแอปเปิ้ลในตะกร้า $>$ 5 ลูก

เมื่อทำการสรุปผลการทดสอบสมมติฐาน :



นางสาวญาญ่าสรุป ยอมรับ H_0 นั่นคือ มีแอปเปิ้ลในตะกร้า $\leq$ 5 ลูก ทั้งที่ความเป็นจริงมีแอปเปิ้ลในตะกร้าจำนวน 6 ลูก กรณีนี้เป็นความผิดพลาดประเภทที่ 2 เนื่องจากผู้วิจัยยอมรับสมมติฐานหลัก (H_0) ทั้งที่ในความเป็นจริงนั้นสมมติฐาน H_0 ไม่จริง

ในการทดสอบสมมติฐานทางสถิติมีโอกาสที่จะตัดสินใจคลาดเคลื่อนหรือผิดพลาดได้ทั้ง 2 ประเภท โดยทั่วไปผู้วิจัยต้องการที่จะทำให้เกิดความผิดพลาดในการทดสอบทั้ง 2 ประเภทน้อยที่สุด แต่อย่างไรก็ตามผู้วิจัยจะไม่สามารถลดความคลาดเคลื่อนทั้ง 2 ประเภทพร้อมกันได้ กล่าวคือ การลด α จะทำให้ β เพิ่มขึ้น ในทำนองเดียวกัน ถ้าลด β จะทำให้ α เพิ่มขึ้น ดังนั้นการที่จะสามารถลดทั้งค่า α และ β จะต้องเพิ่มขนาดตัวอย่าง ในการทดสอบทางสถิตินั้นการที่มีขนาดตัวอย่างมากจะทำให้เกิดความน่าเชื่อถือสูงและถูกต้องมาก นั่นคือ ความคลาดเคลื่อนจะลดน้อยลง

โดยทั่วไปผู้ทดสอบจะกำหนดค่า α (ระดับนัยสำคัญ) หรือกำหนดระดับความเชื่อมั่น $1 - \alpha$ โดยที่ $1 - \alpha$ คือ โอกาสที่จะยอมรับสมมติฐานหลัก (H_0) เมื่อสมมติฐานหลักนั้นเป็นจริง

ความน่าจะเป็นของการปฏิเสธสมมติฐานหลักเมื่อสมมติฐานหลักเป็นเท็จ หรือ $1 - \beta$ จะเรียกว่า อำนาจการทดสอบ หรือ พลังของการทดสอบ (Power of test) เป็นค่าที่บอกให้นักวิจัยทราบถึงความสามารถของการทดสอบสมมติฐานในการจำแนกความแตกต่างที่เกิดขึ้นว่าเป็นความแตกต่างที่แท้จริงหรือไม่ หากมีค่า Power of test สูงแสดงว่าโอกาสที่ผลสรุปจากการทดสอบสมมติฐานจะคลาดเคลื่อนหรือเกิดความผิดพลาดจากความจริงมีน้อย

5.5 เขตวิกฤตและค่าวิกฤต

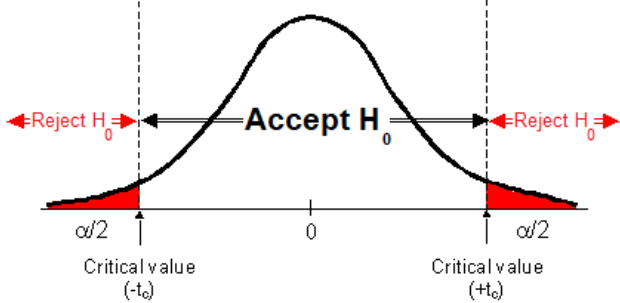
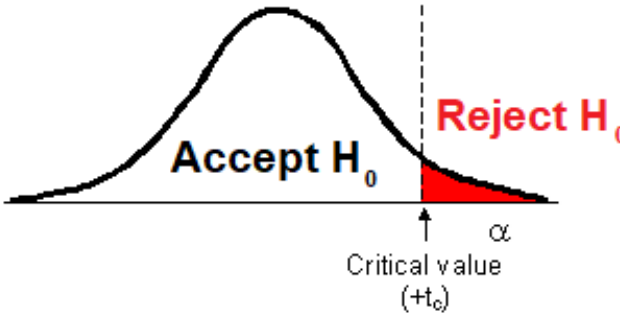
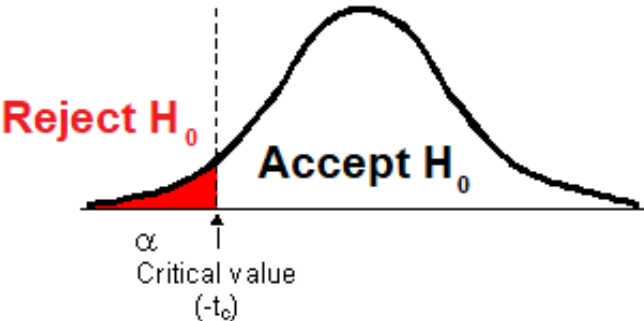
เขตวิกฤต (Critical region) หมายถึง ขอบเขตของการปฏิเสธสมมติฐานหลัก (H_0) ซึ่งกำหนดตามระดับนัยสำคัญ เป็นขอบเขตที่อยู่ทางด้านซ้ายมือ และ/หรือ ขวามือของโค้งการแจกแจงโดยขึ้นกับสมมติฐานรอง (H_1) ดังแสดงในตารางที่ 5.2

ค่าวิกฤต (Critical value) หมายถึง ค่าที่แสดงเขตวิกฤตซึ่งเป็นจุดแบ่งระหว่างเขตยอมรับ (Acceptance region) กับ เขตวิกฤตหรือเขตปฏิเสธสมมติฐาน (Critical region หรือ Rejection region) ดังแสดงในตารางที่ 5.2

เมื่อค่าสถิติอยู่ในเขตวิกฤตหรือเขตปฏิเสธสมมติฐาน (Critical region หรือ Rejection region) จะสรุปได้ว่า ผลการทดสอบมีนัยสำคัญทางสถิติ (Significant) ณ ระดับนัยสำคัญ α ที่กำหนด และมีเหตุผลเพียงพอที่จะสามารถสรุปได้ว่าไม่เป็นไปตามสมมติฐานหลัก (H_0) ที่ตั้งขึ้น จึง ปฏิเสธสมมติฐานหลัก (Reject H_0) และไปยอมรับสมมติฐานรอง (Accept H_1)

แต่ถ้าค่าสถิติตกอยู่ในเขตยอมรับ (Acceptance region) จะสรุปได้ว่า ผลการทดสอบไม่มีนัยสำคัญทางสถิติ (Non-Significant) ณ ระดับนัยสำคัญ α ที่กำหนด และสามารถสรุปได้ว่าไม่มีเหตุผลเพียงพอที่จะสรุปว่าไม่เป็นไปตามสมมติฐานหลัก (H_0) ที่ตั้งขึ้น จึง ยอมรับสมมติฐานหลัก (Accept H_0)

ตารางที่ 5.2 เขตวิกฤตตามประเภทของการทดสอบสมมติฐานทางสถิติ

ประเภทของการทดสอบสมมติฐาน	เขตวิกฤต
การทดสอบแบบสองหาง (Two-tailed test) H_0 : ใช้เครื่องหมาย = H_1 : ใช้เครื่องหมาย $\neq$	
การทดสอบแบบหางเดียว (One-tailed test) ทิศทางมากกว่า H_0 : ใช้เครื่องหมาย $\leq$ H_1 : ใช้เครื่องหมาย $>$	
การทดสอบแบบหางเดียว (One-tailed test) ทิศทางน้อยกว่า H_0 : ใช้เครื่องหมาย $\geq$ H_1 : ใช้เครื่องหมาย $<$	

5.6 ขั้นตอนการทดสอบสมมติฐานทางสถิติ

ขั้นตอนการทดสอบสมมติฐานทางสถิติ เริ่มจากการแปลงสมมติฐานทางการวิจัยให้เป็นสมมติฐานทางสถิติ แล้วใช้เทคนิคทางสถิติแบบใดแบบหนึ่งแล้วแต่กรณีมาทดสอบโดยอาศัยข้อมูลจากกลุ่มตัวอย่าง ขั้นตอนการทดสอบสมมติฐานมีทั้งหมด 5 ขั้นตอน สรุปได้ดังนี้

ขั้นตอนที่ 1 ตั้งสมมติฐานทางสถิติ ผู้วิจัยต้องสามารถกำหนดและเขียนสมมติฐานของงานวิจัยที่คาดไว้หรือเชื่อไว้ให้เป็นสมมติฐานทางสถิติเพื่อการทดสอบ ได้แก่ สมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1)

ขั้นตอนที่ 2 กำหนดระดับนัยสำคัญของการทดสอบ (Level of significance) หรือความผิดพลาดของการทดสอบ เป็นการกำหนดความน่าจะเป็นที่ยอมให้เกิดความคลาดเคลื่อนประเภทที่ 1 (Type I error) ในการทดสอบสมมติฐานทางสถิติ โดยทั่วไปในงานวิจัยด้านพัฒนาผลิตภัณฑ์มักกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 หรือ 0.01

ขั้นตอนที่ 3 คำนวณค่าสถิติทดสอบโดยเลือกเทคนิคทางสถิติที่เหมาะสมเพื่อทดสอบสมมติฐานที่ตั้งไว้ในขั้นตอนที่ 1 คำนวณค่าสถิติทดสอบจากข้อมูลที่นักวิจัยได้รวบรวมไว้ได้จากกลุ่มตัวอย่าง เทคนิคทางสถิติที่นิยมใช้สำหรับการทดสอบสมมติฐานทางสถิติ ได้แก่ z-test, t-test, χ^2 และ F-test

ขั้นตอนที่ 4 ระบุเขตวิกฤตหรือค่าวิกฤต (Critical value) โดยมีหลักในการพิจารณาคือ 1) การทดสอบสมมติฐานเป็นการทดสอบแบบหางเดียว (One-tailed test) หรือสองหาง (Two-tailed test) 2) ค่าระดับนัยสำคัญ (α) ที่กำหนดไว้อยู่ระดับใด เช่น α เท่ากับ 0.05 และ 3) ผู้วิจัยเลือกใช้เทคนิคทางสถิติใดในการทดสอบค่าสถิติ

ขั้นตอนที่ 5 การสรุปผลการทดสอบทางสถิติ นำค่าสถิติที่ได้จากขั้นตอนที่ 3 และขั้นตอนที่ 4 มาเปรียบเทียบกับเพื่อสรุปผล เป็นการสรุปผลที่ได้จากการคำนวณค่าสถิติจากกลุ่มตัวอย่างเพื่ออ้างอิงไปสู่ประชากรที่ศึกษา ถ้าค่าสถิติทดสอบที่คำนวณได้จากข้อ 3 อยู่ในเขตวิกฤตจะสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (Reject H_0) แต่ถ้าหากอยู่นอกเขตวิกฤตจะสรุปได้ว่า ยอมรับสมมติฐานหลัก (Accept H_0)

5.7 หลักเกณฑ์ในการพิจารณาเลือกใช้สถิติในการทดสอบ

การเลือกสถิติเพื่อใช้ในการทดสอบสมมติฐานงานวิจัยขึ้นอยู่กับวัตถุประสงค์ของงานวิจัยและชนิดของข้อมูลที่นำมาทดสอบ ซึ่งสามารถสรุปได้ ดังนี้

5.7.1 การเลือกใช้สถิติในการทดสอบข้อมูลที่มีตัวแปร 1 ตัว (Univariate data analysis)

เป็นการใช้ตัวแปรเพียง 1 ตัวในการวิเคราะห์ข้อมูล การเลือกใช้เทคนิคสถิติจะแตกต่างกันไปขึ้นกับมาตรวัดหรือสเกลในการวัด สามารถแบ่งพิจารณาได้เป็น 2 กรณี คือ

(1) การทดสอบค่าเฉลี่ยของตัวแปรเชิงปริมาณ 1 ตัว

ข้อมูลตัวแปรที่นำมาทดสอบค่าเฉลี่ยต้องเป็นสเกลแบบอันตรภาค (Interval scale) หรือสเกลแบบอัตราส่วน (Ratio scale) เช่น การทดสอบว่าค่าเฉลี่ยคะแนนความชอบมีค่ามากกว่า 6 หรือไม่ ถ้าข้อมูลมีการแจกแจงแบบปกติจะใช้สถิติทดสอบ t-test หรือ z-test แต่ถ้าข้อมูลมีการแจกแจงแบบไม่ปกติจะใช้สถิติแบบนันทพาราเมตริก

การพิจารณาว่าจะใช้สถิติทดสอบ t-test หรือ z-test ขึ้นอยู่กับขนาดตัวอย่าง และการทราบค่าความแปรปรวนของประชากร กล่าวคือ จะใช้สถิติ z-test เมื่อกลุ่มตัวอย่างขนาดใหญ่และทราบค่าความแปรปรวนของประชากร และจะใช้สถิติ t-test เมื่อกลุ่มตัวอย่างมีขนาดเล็ก ($n < 30$) และไม่ทราบค่าความแปรปรวนของประชากร ในทางปฏิบัตินักวิจัยมักไม่ทราบค่าความแปรปรวนของประชากร จึงมักเลือกใช้สถิติ t-test ประกอบกับโปรแกรมสำเร็จรูป SPSS ไม่มีเมนูคำสั่ง z-test มีแต่ t-test และในกรณีกลุ่มตัวอย่างมีขนาดใหญ่ ค่า t-statistics จะมีค่าใกล้เคียงกับ z-statistics ทำให้นักวิจัยสามารถใช้การวิเคราะห์ t-test แทน z-test ได้ในโปรแกรมสำเร็จรูป SPSS

(2) การทดสอบค่าสัดส่วนของตัวแปรเชิงกลุ่ม 1 ตัว

ข้อมูลตัวแปรที่นำมาทดสอบสัดส่วนหรือร้อยละของตัวแปรเชิงกลุ่มต้องเป็นสเกลแบบนามกำหนด (Nominal scale) หรือสเกลอันดับ (Ordinal scale) เช่น ผู้วิจัยคาดว่าผู้บริโภคจะยอมรับผลิตภัณฑ์น้ำอัดลมสูตรลดน้ำตาลในสัดส่วนที่ต่ำกว่า 0.7 (หรือผู้บริโภคจะยอมรับผลิตภัณฑ์น้ำอัดลมสูตรลดน้ำตาลน้อยกว่าร้อยละ 70) สถิติที่ใช้ทดสอบสัดส่วนประชากรของตัวแปรเชิงกลุ่ม 1 ตัว ได้แก่ สถิติทดสอบปกติมาตรฐาน (z-test) และการทดสอบทวินาม (Binomial test) การพิจารณาว่าจะใช้สถิติตัวไหนขึ้นอยู่กับ ขนาดตัวอย่าง และ ค่าสัดส่วนตัวอย่าง

- สถิติทดสอบปกติมาตรฐาน (z-test) ใช้เมื่อตัวอย่างมีขนาดใหญ่ และขนาดตัวอย่างขึ้นกับค่าสัดส่วนตัวอย่างด้วย เช่น ขนาดตัวอย่างมากกว่าหรือเท่ากับ 30 และค่าสัดส่วนตัวอย่าง ($\hat{p}$) เท่ากับ 0.5 สามารถใช้ z-test ทดสอบได้

- การทดสอบทวินาม (Binomial test) ใช้เมื่อตัวอย่างมีขนาดเล็ก ($n < 30$)

ตัวอย่างการเลือกใช้สถิติทดสอบสัดส่วนของประชากร เช่น

กรณีที่ 1 ในการสอบถามนิสิตจำนวน 25 คน เกี่ยวกับความคิดเห็นที่จะให้อาจารย์สอนพิเศษเพิ่มในวันเสาร์และอาทิตย์ พบว่านิสิตสนใจเรียนจำนวน 13 คน และไม่สนใจเรียนจำนวน 12 คน ถ้าต้องการทดสอบว่านิสิตสนใจเรียนพิเศษเท่ากับ 60% หรือไม่ กรณีนี้ต้องใช้การทดสอบทวินาม (Binomial test) เนื่องจากมีขนาดตัวอย่างเล็ก ($n=25$)

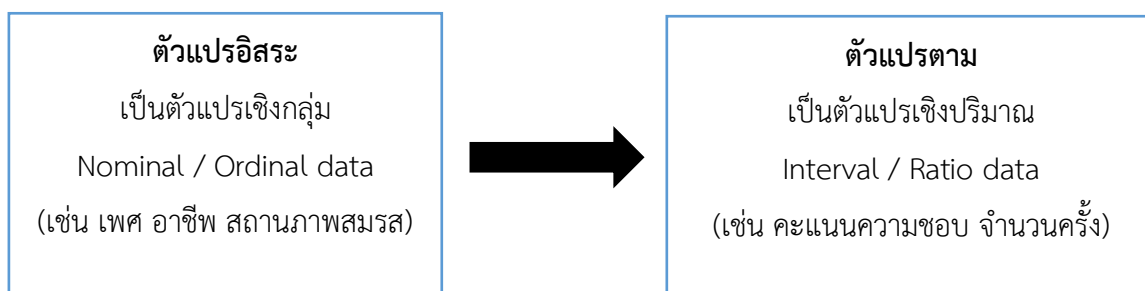
กรณีที่ 2 ในการสอบถามนิสิตม.เกษตร เกี่ยวกับการใส่ชุดนิสิตมาเรียน ผู้วิจัยคาดว่าจะมีนิสิตมากกว่า 60% อยากใส่ชุดนิสิตมาเรียน จึงสุ่มตัวอย่างนิสิตจำนวน 500 คนมาสอบถาม พบว่านิสิตจำนวน 359 คนอยากใส่ชุดนิสิตมาเรียน คิดเป็นค่าสัดส่วนตัวอย่าง $= 359/500 = 0.718$ กรณีนี้ต้องใช้การทดสอบปกติมาตรฐาน (z-test) เนื่องจากมีขนาดตัวอย่างใหญ่ ($n=500$) และค่าสัดส่วนตัวอย่างมากกว่า 0.5

5.7.2 การเลือกใช้สถิติในการทดสอบข้อมูลที่มีตัวแปร 2 ตัว (Bivariate data analysis)

เป็นการวิเคราะห์ข้อมูลตัวแปร 2 ตัวพร้อมกัน ในการวิจัยผู้วิจัยมักต้องการศึกษาหาความสัมพันธ์ระหว่างตัวแปรเพื่อหาเหตุและผล การวิเคราะห์ความสัมพันธ์ของ 2 ตัวแปรเป็นการหาความสัมพันธ์ของตัวแปรอิสระและตัวแปรตามซึ่งอาจจะเป็นชนิดเดียวกันหรือต่างชนิดก็ได้

การเลือกใช้สถิติในการทดสอบข้อมูลที่มีตัวแปร 2 ตัว สามารถพิจารณาออกได้เป็น 3 กรณี ดังนี้

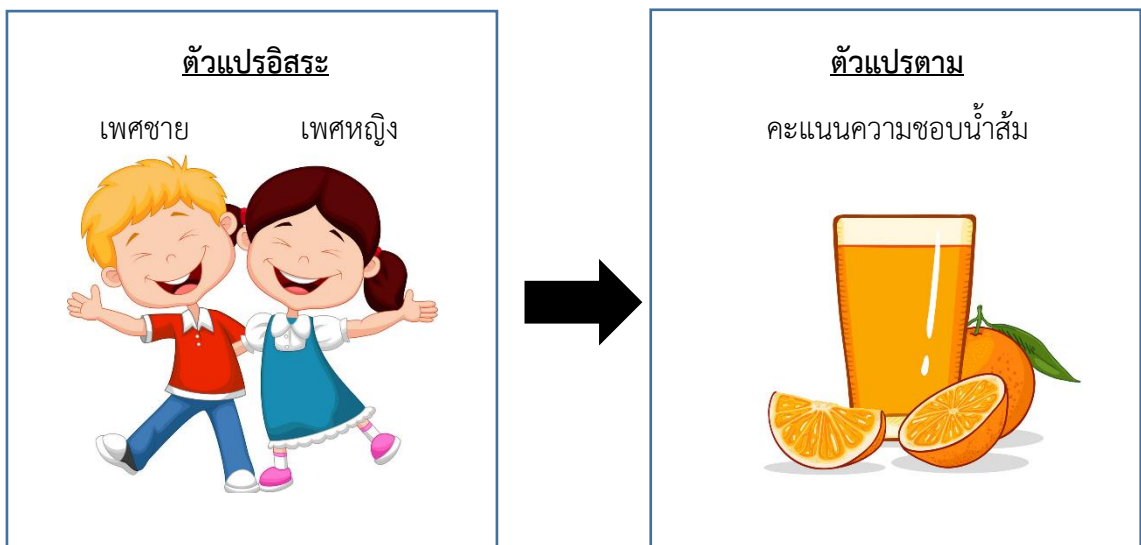
กรณีที่ 1 ตัวแปรอิสระเป็นตัวแปรเชิงกลุ่ม (สเกลนามกำหนด หรือสเกลอันดับ) และตัวแปรตามเป็นตัวแปรเชิงปริมาณ (สเกลอันตรภาค หรือ สเกลอัตราส่วน)



กรณี 1.1 ตัวแปรอิสระเป็นตัวแปรเชิงกลุ่มที่แบ่งเป็น 2 กลุ่มย่อย

ตัวแปรอิสระเป็นตัวแปรเชิงกลุ่มที่แบ่งเป็น 2 กลุ่มย่อย เช่น ตัวแปรเพศ (แบ่งเป็นเพศชาย และเพศหญิง), ตัวแปรอาชีพ (แบ่งเป็นรับราชการ และไม่รับราชการ), ตัวแปรอายุ (แบ่งเป็น น้อยกว่าหรือเท่ากับ 60 ปี และมากกว่า 60 ปี), ตัวแปรสถานภาพนิสิต (แบ่งเป็นนิสิตรอพินิจ และนิสิตปกติ), ตัวแปรการแพ้อาหาร (แบ่งเป็นแพ้อาหาร และไม่แพ้อาหาร), ตัวแปรสถานภาพสมรส (แบ่งเป็นแต่งงานแล้ว และยังไม่ได้แต่งงาน)

สถิติที่ใช้ทดสอบ คือ t-test หรือ z-test



กรณีนักวิจัยคาดว่า เพศชายให้คะแนนความชอบน้ำส้มไม่แตกต่างกับเพศหญิง

ตัวอย่างการตั้งสมมติฐานแบบสองหาง :

H_0 : ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศชาย = ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศหญิง

H_1 : ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศชาย $\neq$ ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศหญิง

กรณีนักวิจัยคาดว่า เพศชายให้คะแนนความชอบน้ำส้มน้อยกว่าเพศหญิง

ตัวอย่างการตั้งสมมติฐานแบบหางเดียว :

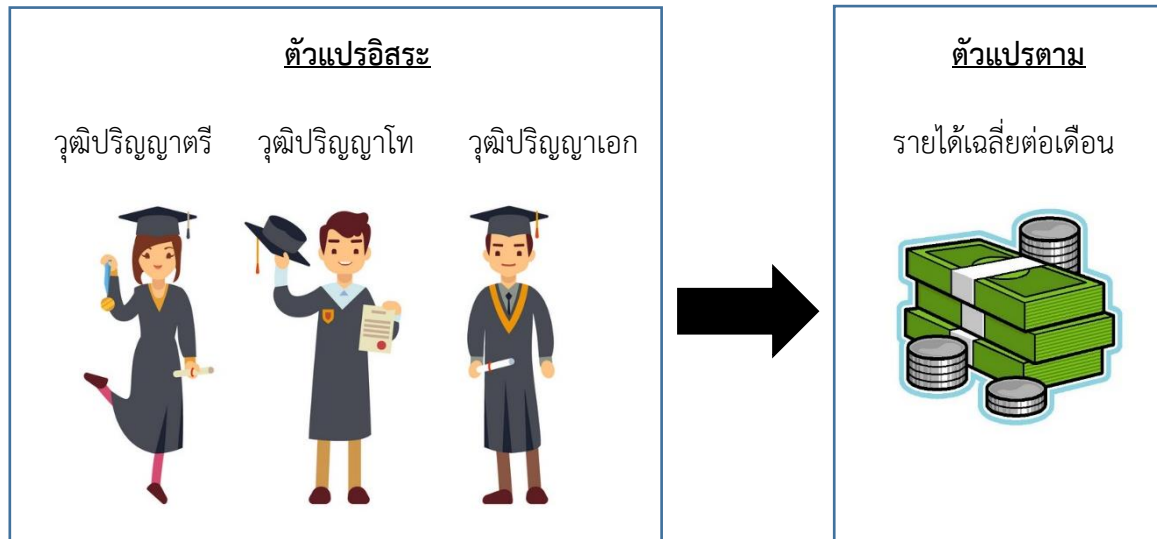
H_0 : ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศชาย $\geq$ ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศหญิง

H_1 : ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศชาย $<$ ค่าเฉลี่ยคะแนนความชอบน้ำส้มของเพศหญิง

กรณี 1.2 ตัวแปรอิสระเป็นตัวแปรเชิงกลุ่มที่แบ่งเป็น 3 กลุ่มย่อยขึ้นไป

ตัวแปรอิสระเป็นตัวแปรเชิงกลุ่มที่แบ่งเป็น 3 กลุ่มย่อยขึ้นไป เช่น ตัวแปรวุฒิการศึกษา (แบ่งเป็นปริญญาตรี, ปริญญาโท และ ปริญญาเอก), ตัวแปรอาชีพ (แบ่งเป็น นิสิต, อาจารย์ และ เจ้าหน้าที่)

สถิติที่ใช้ทดสอบ คือ การวิเคราะห์ความแปรปรวนแบบทางเดียว (One-way ANOVA)



กรณีนี้นักวิจัยคาดว่า วุฒิกศึกษามีผลต่อรายได้เฉลี่ยต่อเดือนของกลุ่มตัวอย่าง

ตัวอย่างการตั้งสมมติฐาน

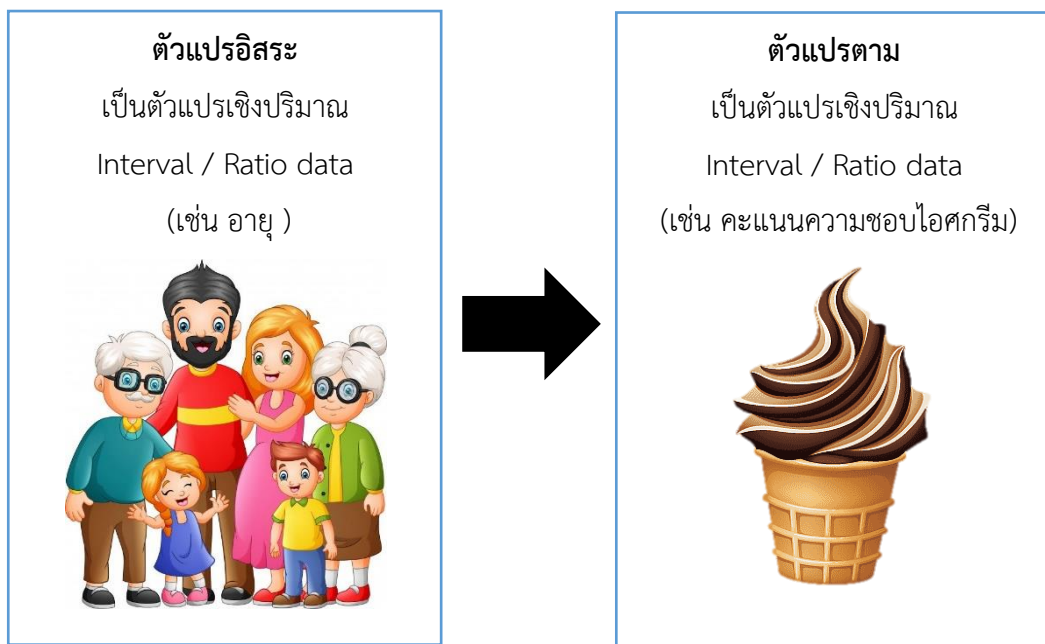
H_0 : ค่าเฉลี่ยรายได้ต่อเดือนของผู้จบการศึกษาทุกระดับวุฒิการศึกษาไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยรายได้ต่อเดือนของผู้จบการศึกษาอย่างน้อย 2 ระดับวุฒิการศึกษาแตกต่างกัน

กรณีที่ 2 ตัวแปรอิสระและตัวแปรตามเป็นตัวแปรเชิงปริมาณ (สเกลอันดับภาค หรือ สเกลอัตราส่วน)

ตัวอย่างเช่น นักวิจัยคาดว่า คະแนนความชอบไอศกรีมช็อกโกแลต (ตัวแปรตาม) ขึ้นอยู่กับอายุของผู้ทดสอบ (ตัวแปรอิสระ) กรณีนี้จะเลือกใช้เทคนิคทางสถิติเพื่อวัดความสัมพันธ์ของทั้งสองตัวแปร

สถิติที่ใช้ทดสอบ คือ การวิเคราะห์ความถดถอยอย่างง่าย (Simple regression) และใช้สัมประสิทธิ์สหสัมพันธ์เพียร์สัน (Pearson correlation coefficient) ในการวัดระดับและทิศทางการสัมพันธ์



กรณีนักวิจัยคาดว่า อายุมีความสัมพันธ์เชิงเส้นตรงกับคະแนนความชอบไอศกรีม

ตัวอย่างการตั้งสมมติฐาน

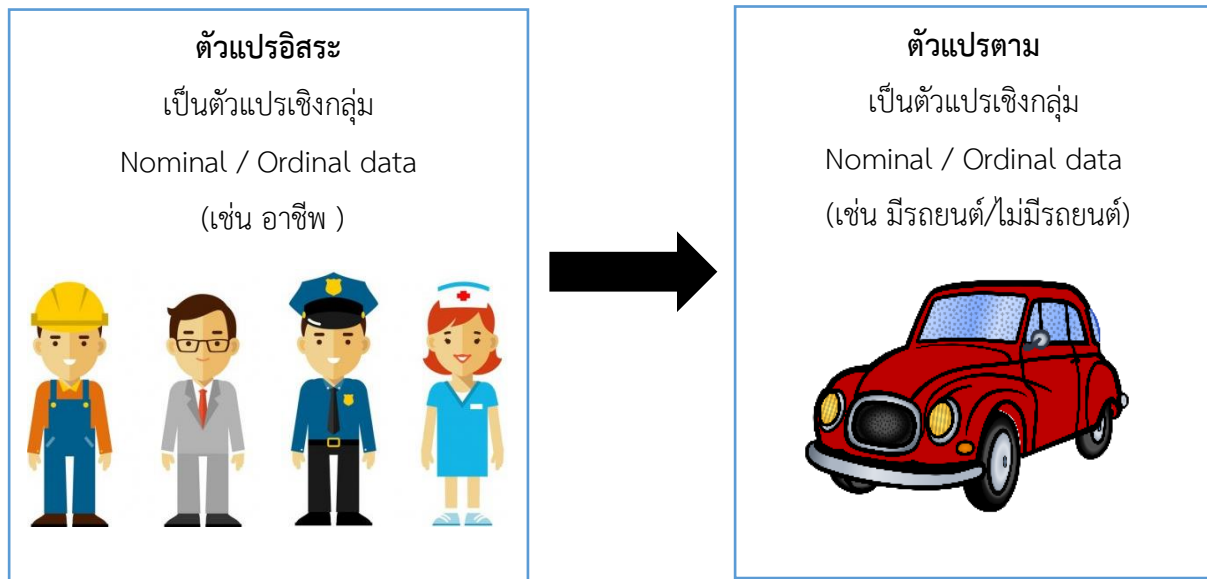
H_0 : คະแนนความชอบไอศกรีมไม่มีความสัมพันธ์กับอายุในรูปแบบเส้นตรง

H_1 : คະแนนความชอบไอศกรีมมีความสัมพันธ์กับอายุในรูปแบบเส้นตรง

กรณีที่ 3 ตัวแปรอิสระและตัวแปรตามเป็นตัวแปรเชิงกลุ่ม (สเกลนามกำหนด หรือ สเกลอันดับ)

ตัวอย่างเช่น นักวิจัยคาดว่า อาชีพ (สเกลนามกำหนด) มีความสัมพันธ์กับการมีรถยนต์ส่วนบุคคล (สเกลนามกำหนด) หรือ นักวิจัยคาดว่า เพศ (สเกลนามกำหนด) มีความสัมพันธ์กับการให้ลำดับความสำคัญของปัจจัยในการเลือกซื้อรถยนต์ (สเกลอันดับ) กรณีนี้จะวัดความสัมพันธ์ระหว่างตัวแปรอิสระและตัวแปรตาม

สถิติที่ใช้ทดสอบ คือ สถิติทดสอบเพียร์สันไคสแควร์ (Pearson Chi-Square)



กรณีนักวิจัยคาดว่า อาชีพมีความสัมพันธ์กับการมี/ไม่มีรถยนต์ส่วนบุคคล

ตัวอย่างการตั้งสมมติฐาน :

H_0 : อาชีพไม่มีความสัมพันธ์กับการมีรถยนต์ส่วนบุคคล

H_1 : อาชีพมีความสัมพันธ์กับการมีรถยนต์ส่วนบุคคล

กรณีที่ 4 ตัวแปรตามเป็นตัวแปรเชิงปริมาณหลายตัว และ ตัวแปรอิสระเป็นตัวแปรเชิงกลุ่มอย่างน้อย 1 ตัว เช่น ผู้วิจัยคาดว่าคะแนนความชอบเค้กและความถี่ในการรับประทาน ขึ้นอยู่กับอายุผู้บริโภค (กำหนดเป็นช่วง) กรณีนี้ตัวแปรอิสระมี 1 ตัว คือ อายุ (ตัวแปรเชิงกลุ่ม แบ่งเป็น 2 กลุ่ม คือ อายุน้อยกว่าเท่ากับ 35 ปี และ อายุมากกว่า 35 ปี) และตัวแปรตามมี 2 ตัว คือ คะแนนความชอบเค้ก และ ความถี่ในการรับประทาน (ตัวแปรเชิงปริมาณ) เทคนิคทางสถิติที่ใช้ทดสอบ คือ MANOVA (Multivariate Analysis of Variance)

ตัวแปรอิสระ



ตัวแปรตาม





บทที่ 6

การวิเคราะห์ผลทางสถิติด้วย t-test

การทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยสำหรับ 1 กลุ่มตัวอย่าง และ 2 กลุ่มตัวอย่าง ใช้เทคนิคทางสถิติที่เรียกว่า t-test ในการทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยสำหรับ 1 กลุ่มตัวอย่าง มีวัตถุประสงค์เพื่อตรวจสอบว่าคุณลักษณะหนึ่งของข้อมูลเป็นไปตามที่คาดหวังหรือกำหนดไว้หรือไม่ โดยพิจารณาจากค่าเฉลี่ย เช่น พนักงานควบคุมคุณภาพของโรงงานต้องการตรวจสอบว่าวัตถุดิบที่เข้ามารอบนี้มีค่าเฉลี่ยปริมาณโปรตีนเท่ากับ 20 เปอร์เซ็นต์ ตามเกณฑ์ของโรงงานหรือไม่ การทดสอบแบบนี้จัดอยู่ในประเภทของการทดสอบแบบ 1 ตัวแปร (Univariate data analysis) ในขณะที่การทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง มีวัตถุประสงค์เพื่อตรวจสอบและเปรียบเทียบว่าคุณลักษณะหนึ่งของข้อมูลระหว่างสองกลุ่มมีความแตกต่างกันหรือไม่ และถ้าแตกต่างกันนั้นแตกต่างอย่างไร โดยพิจารณาจากค่าเฉลี่ยของคุณลักษณะนั้นๆ ความแตกต่างนั้นเป็นความแตกต่างทางสถิติที่ระดับนัยสำคัญตามที่นักวิจัยกำหนด เช่น พนักงานควบคุมคุณภาพของโรงงานต้องการตรวจสอบว่าวัตถุดิบที่เข้ามารอบนี้ของ Supplier A มีปริมาณโปรตีนแตกต่างกับ Supplier B หรือไม่ การทดสอบแบบนี้จัดอยู่ในประเภทของการทดสอบแบบ 2 ตัวแปร (Bivariate data analysis)

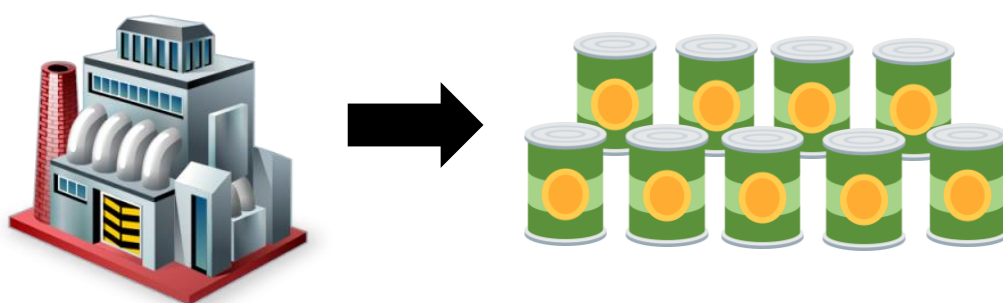
6.1 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 1 กลุ่มตัวอย่าง

วิธี One-sample t-test เป็นการทดสอบสมมติฐานว่าค่าเฉลี่ยของตัวอย่างจำนวน 1 กลุ่มที่มีการแจกแจงแบบปกติหรือใกล้เคียงแบบปกติมีความแตกต่างจากค่าที่ตั้งสมมติฐานไว้หรือไม่ เช่น นิสิตนักศึกษามีค่าใช้จ่ายเฉลี่ยต่อเดือนเท่ากับ 5,000 บาท, มะม่วงอบแห้งมีค่าเฉลี่ยปริมาณความชื้นน้อยกว่า 18 เปอร์เซ็นต์ ทั้งนี้ค่าสมมติฐานที่นักวิจัยตั้งไว้นั้นควรเป็นค่าที่มีเหตุผลสนับสนุนว่าทำไมนักวิจัยจึงคาดเดาไว้เท่านั้น อาจได้มาจากข้อมูลในอดีตที่เคยทำไว้ หรือจากประสบการณ์ของนักวิจัย หรือจากข้อกำหนดมาตรฐานของผลิตภัณฑ์ เช่น มาตรฐานผลิตภัณฑ์ชุมชน (มผช.)

การทดสอบสมมติฐานของค่าเฉลี่ยนั้นข้อมูลที่น่ามาทดสอบจะต้องเป็นข้อมูลเชิงปริมาณที่สามารถนำมาคำนวณได้ ได้แก่ ข้อมูลจากสเกลอันตรภาค (Interval scale) เช่น ระดับความพึงพอใจ และ สเกลอัตราส่วน (Ratio scale) เช่น ปริมาณโปรตีน คคะแนนความชอบ การทดสอบสมมติฐานทำได้ทั้งแบบสองหาง (Two-tailed test) และแบบหางเดียว (One-tailed test) การเขียนสมมติฐานหลัก (H_0) และสมมติฐานรอง (H_1) สามารถเขียนได้ทั้งในรูปแบบข้อความบรรยายหรือใช้สัญลักษณ์ทางคณิตศาสตร์ ดังตัวอย่างต่อไปนี้

6.1.1 การเขียนสมมติฐานทางสถิติเพื่อทดสอบค่าเฉลี่ย 1 กลุ่มตัวอย่าง

ตัวอย่างข้อมูล พนักงานฝ่ายควบคุมคุณภาพของโรงงานแห่งหนึ่งต้องการตรวจสอบว่า ตัวอย่างอาหารกระป๋องที่ผลิตได้มีปริมาณโปรตีนตรงตามที่ระบุในฉลากหรือไม่ (ฉลากระบุว่าปริมาณโปรตีน 20 เปอร์เซ็นต์) จึงได้ทำการสุ่มตัวอย่างอาหารกระป๋องจำนวน 9 กระป๋องที่ผลิตในวันเดียวกันมาวิเคราะห์ปริมาณโปรตีนในห้องปฏิบัติการ ได้ค่าปริมาณโปรตีนแสดงดังตารางที่ 6.1



ตารางที่ 6.1 ปริมาณโปรตีน (%) ของอาหารกระป๋องที่สุ่มตรวจจากโรงงาน

กระป๋องที่	1	2	3	4	5	6	7	8	9
ปริมาณโปรตีน (%)	19.50	20.20	21.50	17.89	18.80	19.80	18.50	18.00	17.90

การตั้งสมมติฐานทางสถิติเพื่อทดสอบ ทำได้ 3 กรณีขึ้นอยู่กับวัตถุประสงค์ของผู้วิจัย ดังนี้

กรณีที่ 1 นักวิจัยต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน เท่ากับ ปริมาณโปรตีน 20 เปอร์เซ็นต์ที่ระบุในฉลากหรือไม่

ประเภทการทดสอบสมมติฐานทางสถิติที่ต้องใช้ คือ การทดสอบแบบสองหาง (Two-tailed test) ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 :	ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน = 20%	หรือ $H_0 : \mu = 20$
H_1 :	ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\neq$ 20%	หรือ $H_1 : \mu \neq 20$

กรณีที่ 2 นักวิจัยต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน มากกว่า ปริมาณโปรตีน 20 เปอร์เซ็นต์ที่ระบุในฉลากหรือไม่

ประเภทการทดสอบสมมติฐานทางสถิติที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) ทิศทางมากกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 :	ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\leq$ 20%	หรือ $H_0 : \mu \leq 20$
H_1 :	ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $>$ 20%	หรือ $H_1 : \mu > 20$

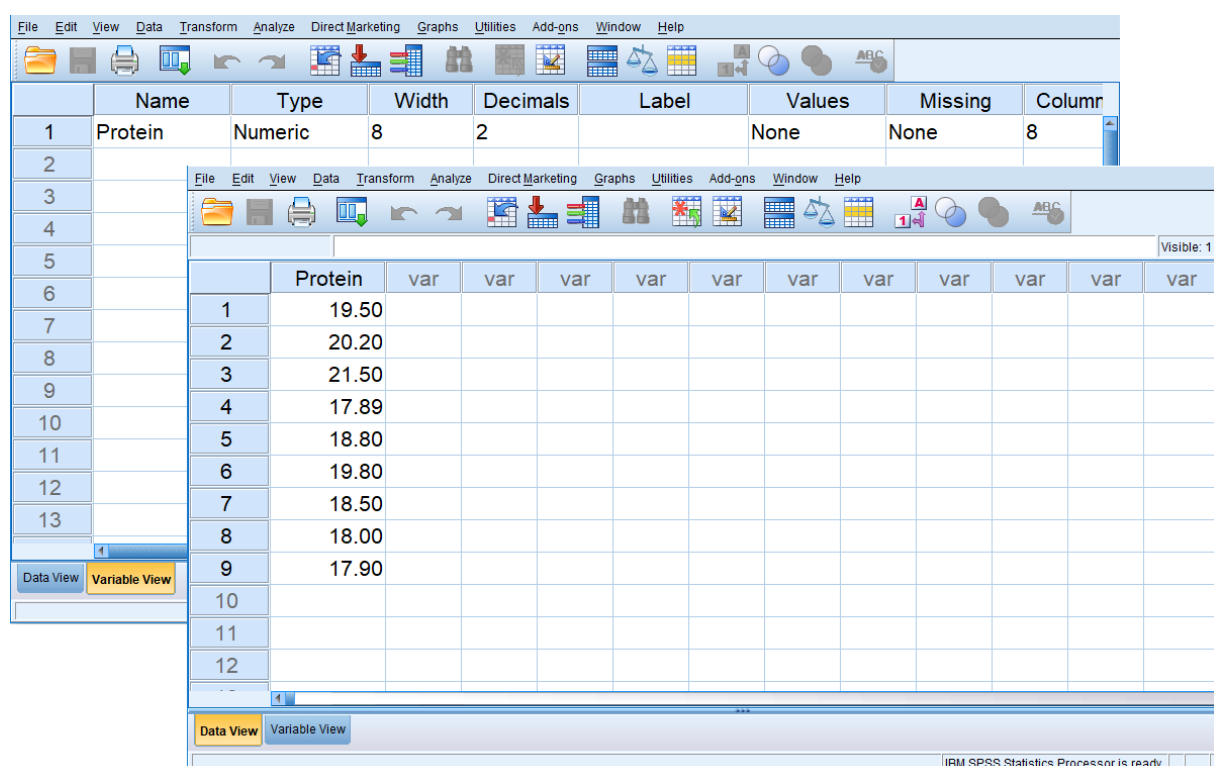
กรณีที่ 3 นักวิจัยต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน น้อยกว่า ปริมาณโปรตีน 20 เปอร์เซ็นต์ที่ระบุในฉลากหรือไม่

ประเภทการทดสอบสมมติฐานทางสถิติที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) ทิศทางน้อยกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 :	ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\geq 20\%$	หรือ $H_0 : \mu \geq 20$
H_1 :	ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $< 20\%$	หรือ $H_1 : \mu < 20$

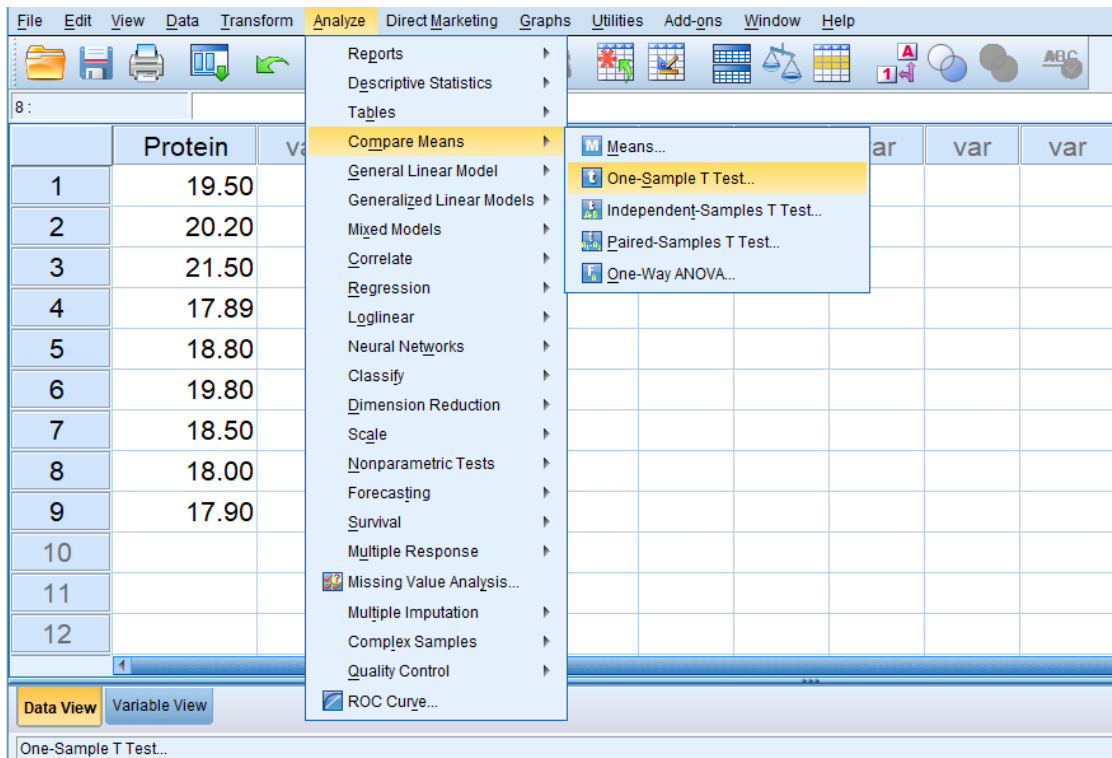
6.1.2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ค่าเฉลี่ยด้วยวิธี One-sample t-test

(1) เปิดโปรแกรม SPSS เพื่อป้อนข้อมูล กำหนดชื่อตัวแปร var เป็น **Protein** ในหน้าต่าง Variable View จากนั้นป้อนข้อมูลปริมาณโปรตีน (จากตารางที่ 6.1) ในหน้าต่าง Data View ดังแสดงในภาพที่ 6.1

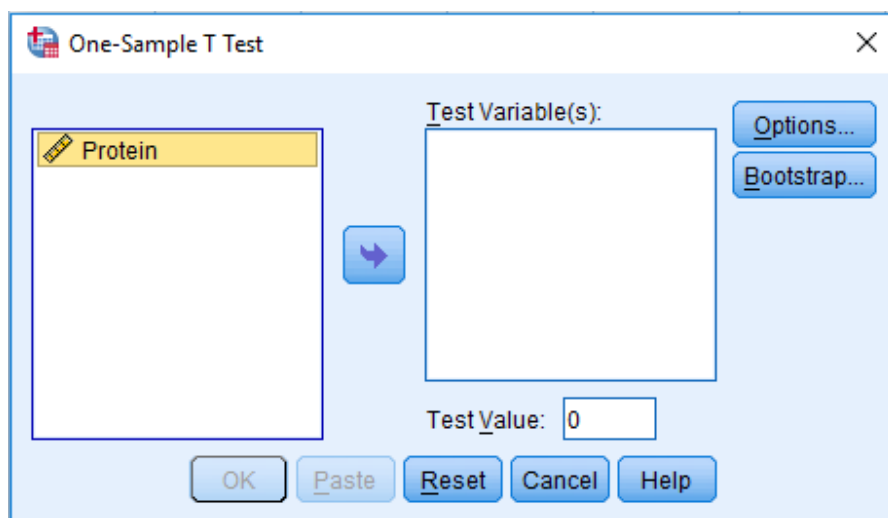


ภาพที่ 6.1 กำหนดชื่อตัวแปร Protein ในหน้าต่าง Variable View และ
ป้อนข้อมูลปริมาณโปรตีนในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > Compare Means > One-Sample T Test...** ดังภาพที่ 6.2 จะปรากฏหน้าจอ ดังภาพที่ 6.3

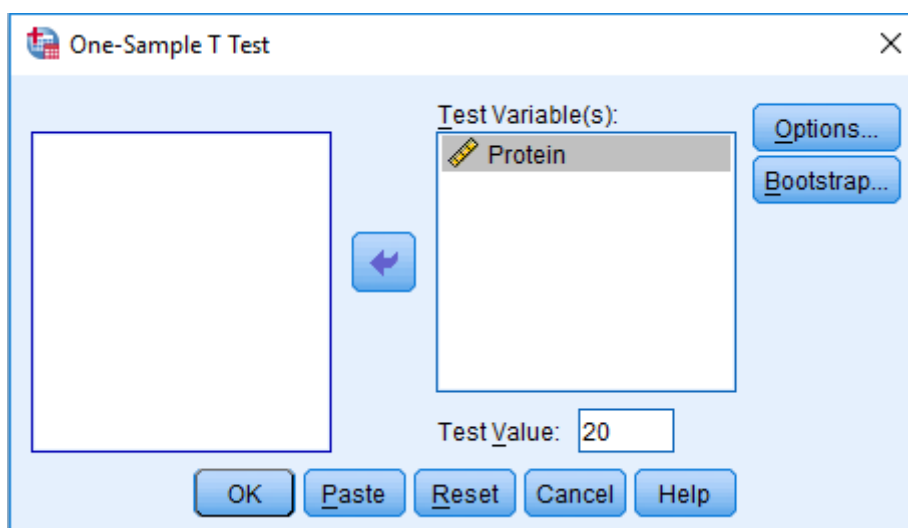


ภาพที่ 6.2 เมนูหน้าจอคำสั่ง Analyze > Compare Means > One-Sample T Test...



ภาพที่ 6.3 หน้าจอคำสั่ง One-Sample T Test

(3) กดคลิกเลือกตัวแปร **Protein** ที่ต้องการวิเคราะห์จากกล่องซ้ายมือไปใส่ใน **Test Variable(s)** ขวามือโดยกดคลิกที่ลูกศร และพิมพ์ค่า **Test Value** เท่ากับ **20** (ซึ่งค่า Test Value คือ ค่า Null hypothesis หรือค่าสมมติฐานที่ตั้งไว้) ดังแสดงในภาพที่ 6.4 แล้วคลิก OK เพื่อวิเคราะห์ผล



ภาพที่ 6.4 การเลือกตัวแปรที่ต้องการวิเคราะห์ และการกำหนดค่าสมมติฐานที่ต้องการทดสอบ

(4) ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 6.5

T-Test

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
Protein	9	19.1211	1.23390	.41130

One-Sample Test

	Test Value = 20					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
Protein	-2.137	8	.065	-.87889	-1.8273	.0696

ภาพที่ 6.5 หน้าจอผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี One-samples t-test

จากผลการวิเคราะห์ในภาพที่ 6.5 แสดงผลลัพธ์ได้ 2 ส่วน ดังนี้

ส่วนที่ 1 ตาราง One-Sample Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

N	คือ จำนวนข้อมูลโปรตีนของอาหารกระป๋องที่นำมาทดสอบ มี 9 ข้อมูล
Mean	คือ ค่าเฉลี่ยปริมาณโปรตีน เท่ากับ 19.1211%
Std. Deviation	คือ ค่าส่วนเบี่ยงเบนมาตรฐานของปริมาณโปรตีน ที่แสดงถึงการกระจายของข้อมูล
Std. Error Mean	คือ ค่าความคลาดเคลื่อนมาตรฐานของข้อมูลปริมาณโปรตีน

ส่วนที่ 2 ตาราง One-Sample Test เป็นส่วนที่แสดงค่าสถิติสำหรับทดสอบสมมติฐานและประมาณค่าโปรตีน จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Test Value	คือ ค่าที่ผู้ทดสอบกำหนดไว้ในสมมติฐาน ในกรณีนี้มีค่าเท่ากับ 20
Mean Difference	คือ ค่าผลต่างของค่าเฉลี่ยของตัวอย่างและค่าที่ผู้ทดสอบกำหนดไว้ในสมมติฐาน (Mean – Test Value) ในที่นี้ค่าที่ได้เป็นลบ แสดงว่า ค่าเฉลี่ยโปรตีนของตัวอย่างที่สุ่มมาน้อยกว่าค่าโปรตีนที่ระบุในฉลาก
95% Confidence Interval of the Difference	คือ ค่าที่แสดงขอบเขตบนล่างของการประมาณค่าผลต่างระหว่างโปรตีนของตัวอย่างและประชากรที่ช่วงความเชื่อมั่น 95%
t, df, Sig.(2-tailed)	คือ ค่าที่โปรแกรมคำนวณได้จากข้อมูลเพื่อใช้ในการตัดสินใจว่าจะยอมรับหรือปฏิเสธสมมติฐานหลัก

หมายเหตุ บางโปรแกรมทางสถิติจะเรียก Sig.(2-tailed) ว่า p-value หรือ Sig.T

6.1.3 การสรุปผลลัพธ์การวิเคราะห์ค่าเฉลี่ยด้วยวิธี One-sample t-test

ข้อมูลผลลัพธ์ที่ได้จากโปรแกรม SPSS นักวิจัยต้องนำมาสรุปผลตามสมมติฐานที่นักวิจัยสนใจและกำหนดไว้ ทั้งนี้การจะยอมรับหรือปฏิเสธสมมติฐานหลัก นักวิจัยต้องพิจารณา 1) ระดับนัยสำคัญที่กำหนดไว้ 2) ประเภทการทดสอบเป็นแบบสองหาง (Two-tailed test) หรือ หางเดียว (One-tailed test) และ 3) พิจารณาค่า t และค่า Sig.(2-tailed) ที่โปรแกรมคำนวณได้ในตาราง One-Sample Test

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 สามารถสรุปเกณฑ์ที่ใช้ในการปฏิเสธสมมติฐานหลัก (H_0) ได้ดังตารางที่ 6.2

ตารางที่ 6.2 สรุปเกณฑ์ที่ใช้ในการปฏิเสธสมมติฐานหลัก (H_0) จากการทดสอบสถิติด้วยวิธี t-test

ประเภทการทดสอบสมมติฐาน	เครื่องหมายของ H_1	ข้อกำหนดในการปฏิเสธ H_0
การทดสอบแบบสองหาง (Two-tailed test)	$\neq$	ค่า $Sig(2 - tailed) \leq \alpha$
การทดสอบแบบหางเดียว (One-tailed test) ทิศทางมากกว่า หรือ ด้านขวา*	$>$	(1) ค่า t เป็นบวก และ (2) ค่า $\frac{Sig.(2-tailed)}{2} \leq \alpha$
การทดสอบแบบหางเดียว (One-tailed test) ทิศทางน้อยกว่า หรือ ด้านซ้าย*	$<$	(1) ค่า t เป็นลบ และ (2) ค่า $\frac{Sig.(2-tailed)}{2} \leq \alpha$

* การทดสอบสมมติฐานแบบหางเดียว ถ้าข้อกำหนดข้อใดข้อหนึ่งไม่จริงจะไม่สามารถปฏิเสธสมมติฐาน H_0 ได้

จากข้อมูลผลลัพธ์ที่ได้ในภาพที่ 6.5 นักวิจัยต้องนำมาสรุปผลตามสมมติฐานที่นักวิจัยสนใจและกำหนดไว้ในข้อ 6.1.1 ในที่นี้จะสรุปผลให้ผู้่านพิจารณาเป็นตัวอย่างทั้ง 3 กรณี ดังนี้

กรณีที่ 1 นักวิจัยต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน เท่ากับ ปริมาณโปรตีน 20 เปอร์เซ็นต์ที่ระบุในฉลากหรือไม่

ประเภทการทดสอบสมมติฐานทางสถิติที่ต้องใช้ คือ การทดสอบแบบสองหาง (Two-tailed test) ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน = 20% หรือ $H_0 : \mu = 20$

H_1 : ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\neq$ 20% หรือ $H_1 : \mu \neq 20$

เนื่องจากการทดสอบแบบสองหาง จะปฏิเสธสมมติฐาน H_0 เมื่อ ค่า $Sig. (2 - tailed) \leq \alpha$

การพิจารณาเพื่อสรุปผล จาก ตาราง One-Sample Test ในภาพที่ 6.5 พบว่า ค่า $Sig.(2-tailed)$ มีค่าเท่ากับ 0.065 ซึ่งมีความมากกว่าค่า α ที่นักวิจัยกำหนดไว้ (ค่า $\alpha = 0.05$)

ดังนั้นจึงตัดสินใจได้ว่า ยอมรับสมมติฐานหลัก (Accept H_0)

สรุปผลได้ว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน = 20% ที่ระดับนัยสำคัญ 0.05

กรณีที่ 2 นักวิจัยต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน มากกว่า ปริมาณโปรตีน 20 เปอร์เซ็นต์ที่ระบุในฉลากหรือไม่

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) ทิศทาง มากกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\leq 20\%$ หรือ $H_0 : \mu \leq 20$

H_1 : ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $> 20\%$ หรือ $H_1 : \mu > 20$

เนื่องจากการทดสอบหางเดียว-ทิศทางมากกว่า จะปฏิเสธสมมติฐาน H_0 ได้เมื่อ (1) ค่า t เป็นบวก และ (2) ค่า $\frac{\text{Sig.}(2\text{-tailed})}{2} \leq \alpha$ ถ้าข้อใดข้อหนึ่งไม่จริงจะไม่สามารถปฏิเสธสมมติฐาน H_0 ได้

การพิจารณาเพื่อสรุปผล จาก ตาราง One-Sample Test ในภาพที่ 6.5 พบว่า มีค่า t เป็นลบ ($t = -2.137$) และค่า Sig.(2-tailed) เมื่อหารสอง มีค่าเท่ากับ 0.0325 ซึ่งมีค่าน้อยกว่าค่า α ที่นักวิจัย กำหนดไว้ (ค่า $\alpha = 0.05$) จึงไม่สามารถปฏิเสธสมมติฐานหลักได้เนื่องจากมีค่า t เป็นลบ

ดังนั้นจึงตัดสินใจได้ว่า ยอมรับสมมติฐานหลัก (Accept H_0)

สรุปผลได้ว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\leq 20\%$ ที่ระดับนัยสำคัญ 0.05

กรณีที่ 3 นักวิจัยต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน น้อยกว่า ปริมาณโปรตีน 20 เปอร์เซ็นต์ที่ระบุในฉลากหรือไม่

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) ทิศทาง น้อยกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $\geq 20\%$ หรือ $H_0 : \mu \geq 20$

H_1 : ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $< 20\%$ หรือ $H_1 : \mu < 20$

เนื่องจากการทดสอบหางเดียว-ทิศทางน้อยกว่า จะปฏิเสธสมมติฐาน H_0 เมื่อ (1) ค่า t เป็นลบ และ (2) ค่า $\frac{\text{Sig.}(2\text{-tailed})}{2} \leq \alpha$ ถ้าข้อใดข้อหนึ่งไม่จริงจะไม่สามารถปฏิเสธสมมติฐาน H_0 ได้

การพิจารณาเพื่อสรุปผล จาก ตาราง One-Sample Test ในภาพที่ 6.5 พบว่า มีค่า t เป็นลบ ($t = -2.137$) และค่า Sig.(2-tailed) เมื่อหารสอง มีค่าเท่ากับ 0.0325 ซึ่งมีค่าน้อยกว่าค่า α ที่นักวิจัย กำหนดไว้ (ค่า $\alpha = 0.05$) จึงสามารถปฏิเสธสมมติฐานหลักได้

ดังนั้นจึงตัดสินใจได้ว่า ปฏิเสธสมมติฐานหลัก (Reject H_0)

สรุปผลได้ว่า ตัวอย่างอาหารกระป๋องมีค่าเฉลี่ยปริมาณโปรตีน $< 20\%$ ที่ระดับนัยสำคัญ 0.05

ข้อสังเกตการสรุปผลลัพธ์จากโปรแกรม SPSS

ค่าระดับนัยสำคัญ (α) ที่นักวิจัยกำหนด ไม่ได้มีผลต่อการคำนวณค่า Sig. (ในตาราง Output) แต่จะมีผลเฉพาะการพิจารณาสรุปผลการทดสอบเท่านั้น กล่าวคือ ไม่ว่านักวิจัยจะกำหนดค่า α เท่ากับ 0.05 หรือ 0.01 ค่า Sig. ที่ได้จากการคำนวณก็ยังคงเท่าเดิม

6.1.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี One-sample t-test

ในการนำเสนอผลลัพธ์ที่ได้จากโปรแกรม SPSS นักวิจัยไม่จำเป็นต้องนำเสนอข้อมูลทุกอย่างที่โปรแกรมให้ผลลัพธ์ออกมา แต่ควรเลือกนำเสนอเฉพาะที่จำเป็นสำหรับการทดสอบที่สนใจ ในการทดสอบสมมติฐานค่าเฉลี่ยของตัวอย่าง 1 กลุ่มนั้นควรนำเสนออย่างน้อย 3 ค่าที่เกี่ยวข้อง ดังนี้ 1) ค่าสถิติที่ใช้ในการทดสอบ, 2) ค่าระดับนัยสำคัญที่ใช้ในการทดสอบ และ 3) ผลสรุปที่ได้จากการทดสอบสมมติฐาน

จากตัวอย่างที่พนักงานฝ่ายควบคุมคุณภาพของโรงงานต้องการทดสอบว่า ตัวอย่างอาหารกระป๋องที่สุ่มมาตรวจมีค่าเฉลี่ยปริมาณโปรตีน เท่ากับ ปริมาณโปรตีน 20 เปอร์เซ็นต์ตามที่ระบุในฉลากหรือไม่ และได้ผลลัพธ์การวิเคราะห์ทางสถิติด้วยวิธี One-sample t-test ดังภาพที่ 6.5 นั้น สามารถนำเสนอผลการวิเคราะห์ทางสถิติในรูปแบบตารางได้ดังตารางที่ 6.3

ตารางที่ 6.3 ผลการเปรียบเทียบค่าเฉลี่ยปริมาณโปรตีนของอาหารกระป๋องที่สุ่มตรวจกับค่าปริมาณโปรตีนที่ระบุในฉลาก (20 เปอร์เซ็นต์)

ปริมาณโปรตีน (%)	Mean	S.D.	n	t-value	Sig.(2-tailed)
	19.12	1.23	9	-2.137	0.065

กรณีสมมติ หากผลการทดสอบสมมติฐานทางสถิติ สรุปได้ว่า ปฏิเสธสมมติฐานหลัก (Reject H_0) สมมติว่าได้ค่า Sig.(2-tailed) = 0.026 นักวิจัยต้องกำหนดสัญลักษณ์แสดงการปฏิเสธ H_0 ในตาราง โดยใส่เครื่องหมาย * (ดอกจัน) หลังค่า Sig.(2-tailed) และระบุคำอธิบายได้ตารางดังแสดงในตารางที่ 6.4

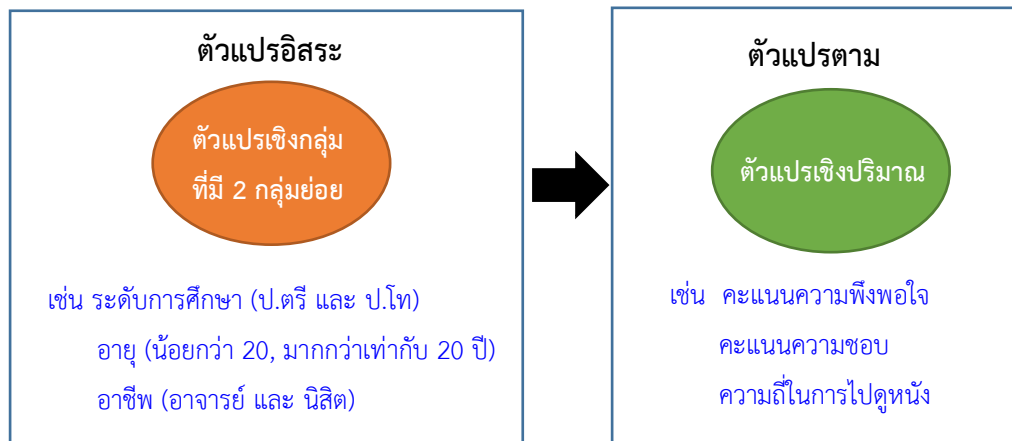
ตารางที่ 6.4 ผลการเปรียบเทียบค่าเฉลี่ยปริมาณโปรตีนของอาหารกระป๋องที่สุ่มตรวจกับค่าปริมาณโปรตีนที่ระบุในฉลาก (20 เปอร์เซ็นต์)

ปริมาณโปรตีน (%)	Mean	S.D.	n	t-value	Sig.(2-tailed)
	19.12	1.23	9	-2.137	0.026*

* ระดับนัยสำคัญ $\alpha = 0.05$

6.2 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง

การเปรียบเทียบค่าเฉลี่ยระหว่าง 2 กลุ่มประชากรนั้น ข้อมูลจะประกอบด้วย (1) ตัวแปรอิสระหรือตัวแปรต้นที่เป็นตัวแปรเชิงกลุ่ม (สเกลนามกำหนด) โดยมี 2 กลุ่มย่อย และ (2) ตัวแปรตามเป็นตัวแปรเชิงปริมาณ (สเกลอันดับภาค หรือสเกลอัตราส่วน) การเปรียบเทียบค่าเฉลี่ยระหว่าง 2 กลุ่มประชากรนั้นอาจกล่าวได้ว่าเป็นการศึกษาความสัมพันธ์ของตัวแปรอิสระและตัวแปรตาม



การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง มีวัตถุประสงค์เพื่อตรวจสอบและเปรียบเทียบว่าคุณลักษณะหนึ่งของข้อมูลระหว่างสองกลุ่มมีความแตกต่างกันหรือไม่ และถ้าแตกต่างกันนั้นแตกต่างอย่างไร โดยพิจารณาจากค่าเฉลี่ยของคุณลักษณะนั้นๆ การใช้สถิติ t-test เปรียบเทียบค่าเฉลี่ยระหว่าง 2 กลุ่มประชากรนั้น มีเงื่อนไขการทดสอบ คือ ประชากรมีการแจกแจงแบบปกติหรือใกล้เคียงปกติ ไม่ทราบค่าแปรปรวนของแต่ละประชากร และกลุ่มตัวอย่างมีขนาดเล็ก

6.2.1 สถิติ t-test ที่ใช้ทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง

แบ่งออกเป็น 2 ประเภทตามความเป็นอิสระของกลุ่มประชากรที่สุ่ม คือ

(1) **สุ่มตัวอย่างจากประชากรที่ไม่เป็นอิสระกัน (Dependent sample)** เป็นการทดสอบความแตกต่างของค่าเฉลี่ยระหว่าง 2 กลุ่มตัวอย่าง เมื่อข้อมูลตัวอย่างที่จะใช้ทดสอบมีความสัมพันธ์กัน เช่น การเปรียบเทียบปริมาณน้ำตาลในเลือดของคนไข้ก่อนและหลังรับประทานยาเป็นเวลา 1 เดือน ซึ่งการทดสอบแบบนี้จะเป็นการทดสอบเป็นคู่ ๆ โดยแต่ละคู่มีความสัมพันธ์กัน จึงอาจเรียกการทดสอบแบบนี้ว่าการทดสอบความแตกต่างแบบจับคู่ (Paired difference test) เทคนิคทางสถิติที่ใช้ในการทดสอบสมมติฐาน คือ Paired-samples t-test

(2) **สุ่มตัวอย่างจากประชากรที่เป็นอิสระกัน (Independent sample)** เป็นการทดสอบความแตกต่างของค่าเฉลี่ยระหว่าง 2 กลุ่มตัวอย่าง เมื่อข้อมูลตัวอย่างที่จะใช้ทดสอบไม่มีความสัมพันธ์กัน เช่น ทดสอบความแตกต่างด้านความหวานของมะม่วงน้ำดอกไม้สวนนาย A และสวนนาย B เทคนิคทางสถิติที่ใช้ในการทดสอบสมมติฐาน คือ Independent-samples t-test

6.2.2 การเขียนสมมติฐานทางสถิติสำหรับทดสอบค่าเฉลี่ย 2 กลุ่มตัวอย่าง

ตัวอย่างที่ 1 นักพัฒนาผลิตภัณฑ์ต้องการศึกษาว่า ผู้ทดสอบให้คะแนนความชอบไอศกรีมสูตรที่ใช้ไซรัปกล้วยแตกต่างกับไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วยหรือไม่ จึงให้ผู้ทดสอบที่เป็นนิสิตระดับปริญญาตรีจำนวน 50 คน ชิมตัวอย่างไอศกรีมทั้ง 2 ตัวอย่าง (ไอศกรีมที่ใช้ไซรัปกล้วย และ ไอศกรีมที่ใช้สารแต่งกลิ่นกล้วย) แล้วประเมินความชอบด้วยสเกล 9-point hedonic scale (โดยคะแนน 1 = ไม่ชอบมากที่สุด และคะแนน 9 = ชอบมากที่สุด)

จากข้อมูลดังกล่าว สามารถพิจารณารายละเอียดของข้อมูลได้ ดังนี้

- (1) ตัวแปรอิสระ คือ สูตรไอศกรีม (มี 2 สูตร ได้แก่ สูตรใช้ไซรัปกล้วย และสูตรใช้สารแต่งกลิ่นกล้วย)
- (2) ตัวแปรตาม คือ ค่าคะแนนความชอบ
- (3) ข้อมูลแต่ละกลุ่มตัวอย่าง คือ คะแนนความชอบที่มาจากกลุ่มผู้ทดสอบเดียวกัน (นิสิตระดับปริญญาตรี จำนวน 50 คน กลุ่มเดียว) ดังนั้น ข้อมูลค่าคะแนนความชอบไอศกรีมทั้งสองสูตรมีความสัมพันธ์กัน
- (4) เทคนิคทางสถิติที่ใช้ในการทดสอบสมมติฐาน คือ Paired-samples t-test

การตั้งสมมติฐานทางสถิติเพื่อการทดสอบ ทำได้ 3 กรณีขึ้นอยู่กับวัตถุประสงค์ผู้วิจัย ดังนี้

กรณีที่ 1 ถ้านักพัฒนาผลิตภัณฑ์สนใจเปรียบเทียบ ความแตกต่าง ระหว่างค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซรัปกล้วยกับไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบสองหาง (Two-tailed test)

ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซรัปกล้วย = ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย

H_1 : ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซรัปกล้วย $\neq$ ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

$$H_0 : \mu_{\text{ไซร้ปกล้วย}} = \mu_{\text{สารแต่งกลิ่นกล้วย}}$$

$$H_1 : \mu_{\text{ไซร้ปกล้วย}} \neq \mu_{\text{สารแต่งกลิ่นกล้วย}}$$

กรณีที่ 2 ถ้านักพัฒนาผลิตภัณฑ์คาดว่า ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซร้ปกล้วย มากกว่า ไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) – ทิศทางมากกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

$$H_0 : \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซร้ปกล้วย} \leq \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย}$$

$$H_1 : \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซร้ปกล้วย} > \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย}$$

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

$$H_0 : \mu_{\text{ไซร้ปกล้วย}} \leq \mu_{\text{สารแต่งกลิ่นกล้วย}}$$

$$H_1 : \mu_{\text{ไซร้ปกล้วย}} > \mu_{\text{สารแต่งกลิ่นกล้วย}}$$

กรณีที่ 3 ถ้านักพัฒนาผลิตภัณฑ์คาดว่า ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซร้ปกล้วย น้อยกว่า ไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) - ทิศทางน้อยกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

$$H_0 : \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซร้ปกล้วย} \geq \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย}$$

$$H_1 : \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้ไซร้ปกล้วย} < \text{ค่าเฉลี่ยคะแนนความชอบไอศกรีมสูตรที่ใช้สารแต่งกลิ่นกล้วย}$$

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

$$H_0 : \mu_{\text{ไซร้ปกล้วย}} \geq \mu_{\text{สารแต่งกลิ่นกล้วย}}$$

$$H_1 : \mu_{\text{ไซร้ปกล้วย}} < \mu_{\text{สารแต่งกลิ่นกล้วย}}$$

ตัวอย่างที่ 2 นักพัฒนาผลิตภัณฑ์ต้องการศึกษาว่า ลูกค้าชาวญี่ปุ่นและลูกค้าชาวเกาหลีใต้ให้คะแนนความชอบเครื่องดื่มชาข้าวไรซ์เบอร์รี่แตกต่างกันหรือไม่ จึงให้ผู้ทดสอบชาวญี่ปุ่นจำนวน 50 คน และชาวเกาหลีใต้จำนวน 50 คน ชิมตัวอย่างเครื่องดื่มชาข้าวไรซ์เบอร์รี่แล้วประเมินความชอบด้วยสเกล 9-point hedonic scale (โดยคะแนน 1 = ไม่ชอบมากที่สุด และคะแนน 9 = ชอบมากที่สุด)

จากข้อมูลดังกล่าว สามารถพิจารณารายละเอียดของข้อมูลได้ ดังนี้

- (1) ตัวแปรอิสระ คือ กลุ่มลูกค้า (มี 2 กลุ่ม ได้แก่ ลูกค้าชาวญี่ปุ่น และลูกค้าชาวเกาหลีใต้)
- (2) ตัวแปรตาม คือ ค่าคะแนนความชอบ
- (3) ข้อมูลแต่ละกลุ่มตัวอย่าง คือ คะแนนความชอบที่มาจากกลุ่มผู้ทดสอบที่แตกต่างกัน (ลูกค้าชาวญี่ปุ่น และลูกค้าชาวเกาหลีใต้) ดังนั้น ข้อมูลทั้งสองกลุ่มไม่มีความสัมพันธ์กัน
- (4) เทคนิคทางสถิติที่ใช้ในการทดสอบสมมติฐาน คือ Independent-samples t-test

การตั้งสมมติฐานทางสถิติเพื่อการทดสอบ ทำได้ 3 กรณีขึ้นอยู่กับวัตถุประสงค์ผู้วิจัย ดังนี้

กรณีที่ 1 ถ้านักพัฒนาผลิตภัณฑ์สนใจเปรียบเทียบ ความแตกต่าง ระหว่างค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่นและลูกค้าชาวเกาหลีใต้

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบสองหาง (Two-tailed test)

ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น = ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

H_1 : ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น $\neq$ ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

H_0 : $\mu_{ญี่ปุ่น} = \mu_{เกาหลีใต้}$

H_1 : $\mu_{ญี่ปุ่น} \neq \mu_{เกาหลีใต้}$

กรณีที่ 2 ถ้านักพัฒนาผลิตภัณฑ์คาดว่า ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น มากกว่า ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) - ทิศทางมากกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น $\leq$ ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

H_1 : ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น $>$ ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

H_0 : $\mu_{\text{ญี่ปุ่น}} \leq \mu_{\text{เกาหลีใต้}}$

H_1 : $\mu_{\text{ญี่ปุ่น}} > \mu_{\text{เกาหลีใต้}}$

กรณีที่ 3 ถ้านักพัฒนาผลิตภัณฑ์คาดว่า ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น น้อยกว่า ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

ประเภทการทดสอบสมมติฐานที่ต้องใช้ คือ การทดสอบแบบหางเดียว (One-tailed test) - ทิศทางน้อยกว่า ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น $\geq$ ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

H_1 : ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวญี่ปุ่น $<$ ค่าเฉลี่ยคะแนนความชอบของลูกค้าชาวเกาหลีใต้

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ ได้ดังนี้

H_0 : $\mu_{\text{ญี่ปุ่น}} \geq \mu_{\text{เกาหลีใต้}}$

H_1 : $\mu_{\text{ญี่ปุ่น}} < \mu_{\text{เกาหลีใต้}}$

6.3 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่างที่ไม่เป็นอิสระกัน

วิธี Paired-samples t-test เป็นการทดสอบค่าเฉลี่ยของ 2 กลุ่มตัวอย่างว่ามีความแตกต่างกันหรือไม่ โดยที่จะทำการสุ่มตัวอย่างทีละคู่ หรืออาจกล่าวได้ว่าการทดสอบจะทดสอบว่าความแตกต่างของแต่ละคู่มีความแตกต่างจากศูนย์หรือไม่ การทดสอบสมมติฐานทางสถิติแบบนี้อาจทำได้ทั้งแบบทางเดียว (One-tailed test) และแบบสองทาง (Two-tailed test) ข้อมูลที่ใช้การวิเคราะห์ด้วยวิธีนี้ต้องเป็นข้อมูลที่ไม่เป็นอิสระกัน หรือ มีความสัมพันธ์กัน จำนวนข้อมูลของแต่ละกลุ่มตัวอย่างจำเป็นต้องเท่ากัน ลักษณะตัวอย่างข้อมูลประเภทนี้ เช่น ผู้ทดสอบคนเดียวกันได้รับการทดสอบชม 2 ตัวอย่าง, คะแนนสอบของผู้เข้ารับการอบรมก่อนและหลังการอบรม, ปริมาณคอเลสเตอรอลของคนไข้ก่อนและหลังใช้ยา

6.3.1 ลักษณะของข้อมูลกลุ่มตัวอย่างที่ไม่เป็นอิสระกัน พิจารณาได้ ดังนี้

(1) การเปรียบเทียบวิธีการ 2 วิธีกับข้อมูลชุดเดียวกัน

เช่น ฝ่ายประกันคุณภาพของบริษัทต้องการเปรียบเทียบความถูกต้องของเครื่องวิเคราะห์ไขมัน 2 ยี่ห้อ (ยี่ห้อ A และยี่ห้อ B) จึงนำตัวอย่างอาหารมา 10 ตัวอย่าง แต่ละตัวอย่างวิเคราะห์ด้วยเครื่องยี่ห้อ A และยี่ห้อ B จากนั้นจึงเปรียบเทียบค่าความแตกต่างของปริมาณไขมันที่วิเคราะห์ได้จากทั้งสองเครื่องมือ

(2) การเปรียบเทียบข้อมูล 2 ชุดกับคุณสมบัติที่เหมือนกัน

เป็นการเปรียบเทียบกลุ่มตัวอย่างที่มีคุณสมบัติเหมือนกันมาจับคู่กัน เช่น นำนิสิตหมู่ 1 และหมู่ 2 ที่มีคะแนนเกรดเฉลี่ยเท่ากันและเรียนวิชาเคมีโดยอาจารย์ผู้สอนคนเดียวกัน มาเปรียบเทียบคะแนนการสอบปลายภาควิชาเคมี โดยทำการเปรียบเทียบเป็นคู่ ๆ

(3) การเปรียบเทียบข้อมูล 2 ประเภทที่ได้จากแหล่งข้อมูลเดียวกัน

เช่น การเปรียบเทียบคะแนนความชอบน้ำส้ม 2 สูตรโดยให้ผู้ทดสอบ 50 คนชิมน้ำส้มทั้งสองสูตรแล้วประเมินความชอบ

(4) การเปรียบเทียบข้อมูล 2 ประเภทที่ได้จากช่วงเวลาเดียวกัน

เช่น บริษัทผลิตนมอัดเม็ดแห่งหนึ่งมีโรงงานผลิต 2 แห่ง (โรงงาน A และโรงงาน B) ในแต่ละวันพบว่ามินมอัดเม็ดที่ต้องคัดทิ้งเนื่องจากการแตกหัก จึงได้ทำการเก็บข้อมูลเพื่อเปรียบเทียบปริมาณนมอัดเม็ดที่คัดทิ้งของทั้งสองโรงงานในแต่ละวันและเวลาเดียวกัน

6.3.2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-Samples t-test

ตัวอย่างข้อมูล นักวิจัยของโรงพยาบาลแห่งหนึ่งต้องการศึกษาผลของการใช้ยาลดปริมาณคอเลสเตอรอลที่มีต่อคนไข้จำนวน 8 ราย เพื่อดูประสิทธิภาพของยา โดยทำการวัดปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยา และหลังใช้ยาเป็นเวลา 1 เดือน ได้ผลการวิเคราะห์แสดงดังตารางที่ 6.5

ตารางที่ 6.5 ปริมาณคอเลสเตอรอลก่อนใช้ยาและหลังใช้ยาของผู้ป่วย 8 ราย

คนไข้คนที่	1	2	3	4	5	6	7	8
ก่อนใช้ยา (mg/dL)	131	120	118	117	109	118	110	147
หลังใช้ยา (mg/dL)	128	117	104	101	100	117	105	136

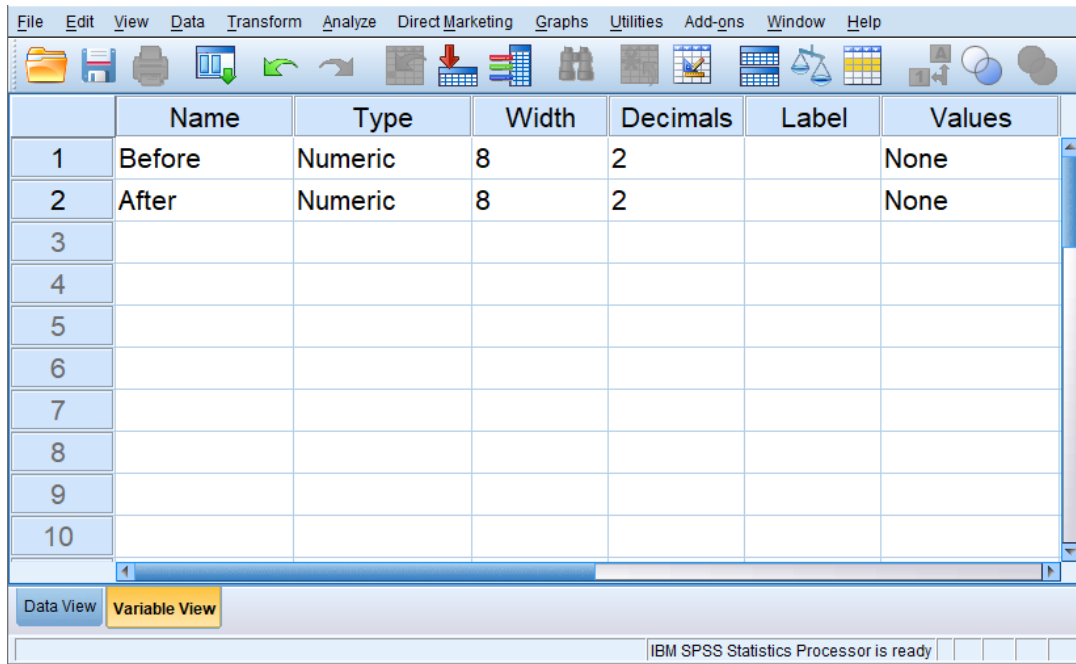
ถ้านักวิจัยของโรงพยาบาลต้องการการเปรียบเทียบว่าปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยาและหลังใช้ยาแตกต่างกันหรือไม่ กรณีนี้ประเภทการทดสอบสมมติฐานทางสถิติที่ต้องใช้ คือ การทดสอบแบบสองหาง (Two-tailed test) ซึ่งสามารถเขียนสมมติฐานหลัก (H_0) และ สมมติฐานรอง (H_1) ได้ดังนี้

H_0 : ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยา = ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้หลังใช้ยา

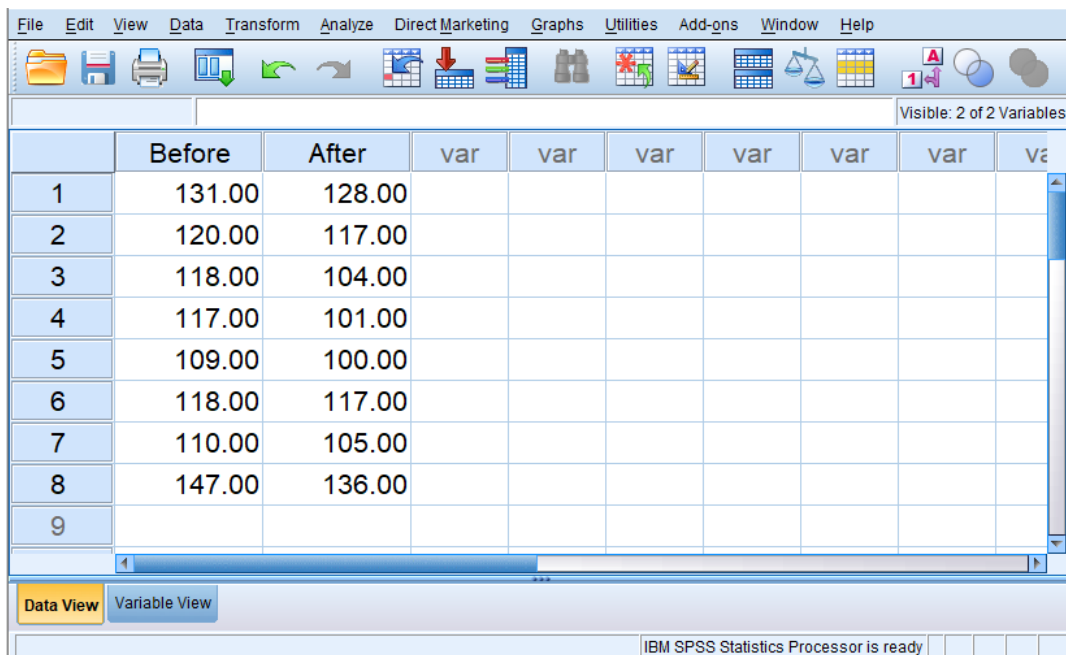
H_1 : ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยา $\neq$ ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้หลังใช้ยา

ขั้นตอนการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test โดยโปรแกรม SPSS

(1) เปิดโปรแกรม SPSS เพื่อป้อนข้อมูล ในกรณีนี้จะมีตัวแปร 2 ตัว คือ ตัวแปรก่อนใช้ยา และตัวแปรหลังใช้ยา ดังนั้นกำหนดชื่อตัวแปรเป็น **Before** และ **After** ในหน้าต่าง Variable View ดังภาพที่ 6.6 จากนั้นป้อนข้อมูลปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยาและหลังใช้ยาเป็นคู่ ๆ ในหน้าต่าง Data View ดังภาพที่ 6.7



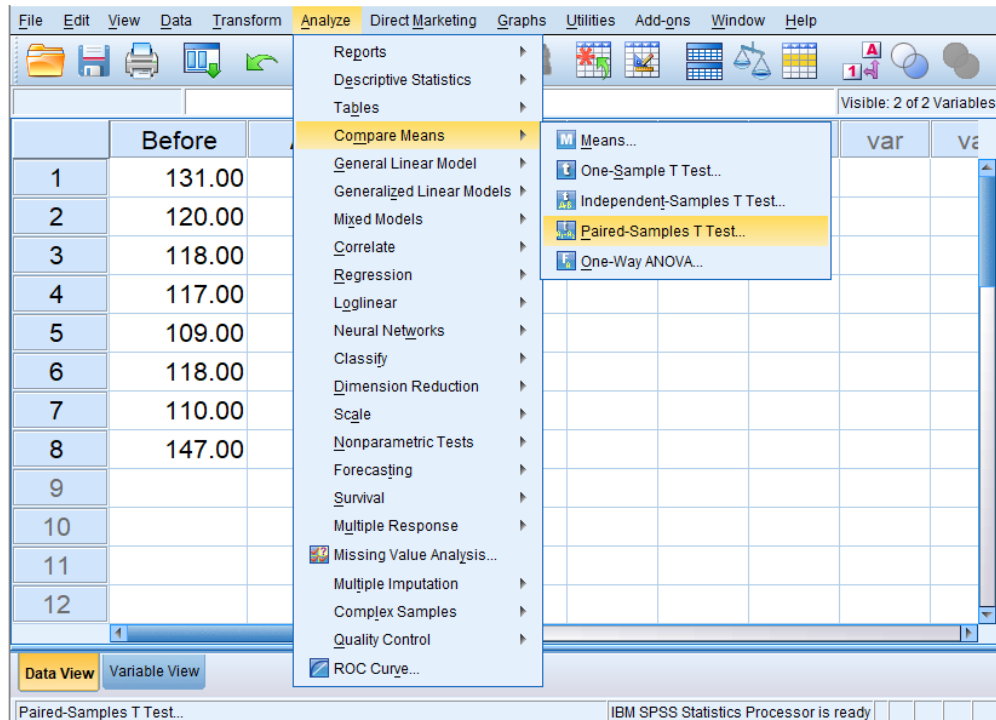
ภาพที่ 6.6 กำหนดชื่อตัวแปร Before และ After ในหน้าต่าง Variable View



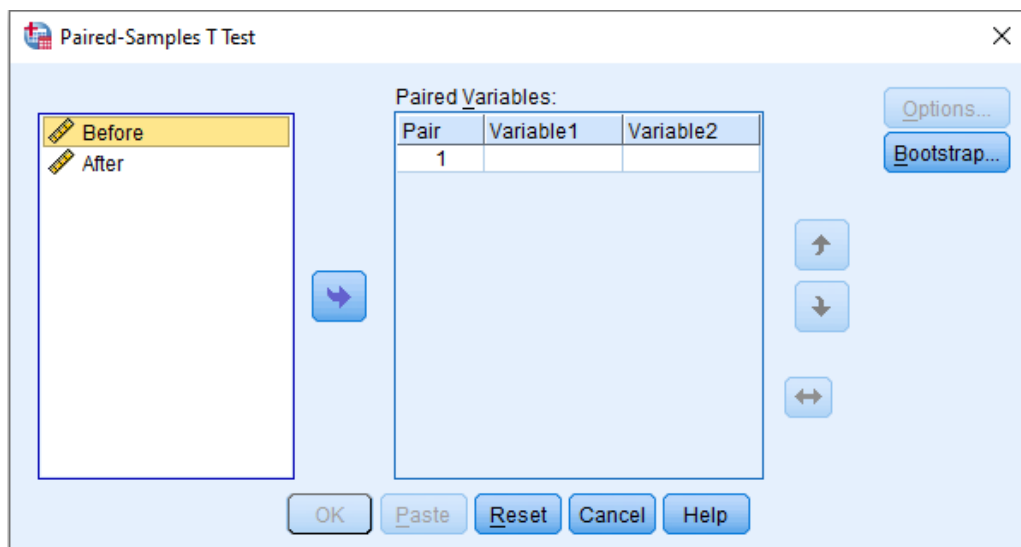
ภาพที่ 6.7 ป้อนข้อมูลคอเลสเตอรอลของคนไข้แต่ละคนเป็นคู่ ๆ (ก่อน-หลังใช้ยา) ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > Compare Means > Paired-Samples T Test...** ดังภาพที่ 6.8

จะปรากฏหน้าต่างดังภาพที่ 6.9

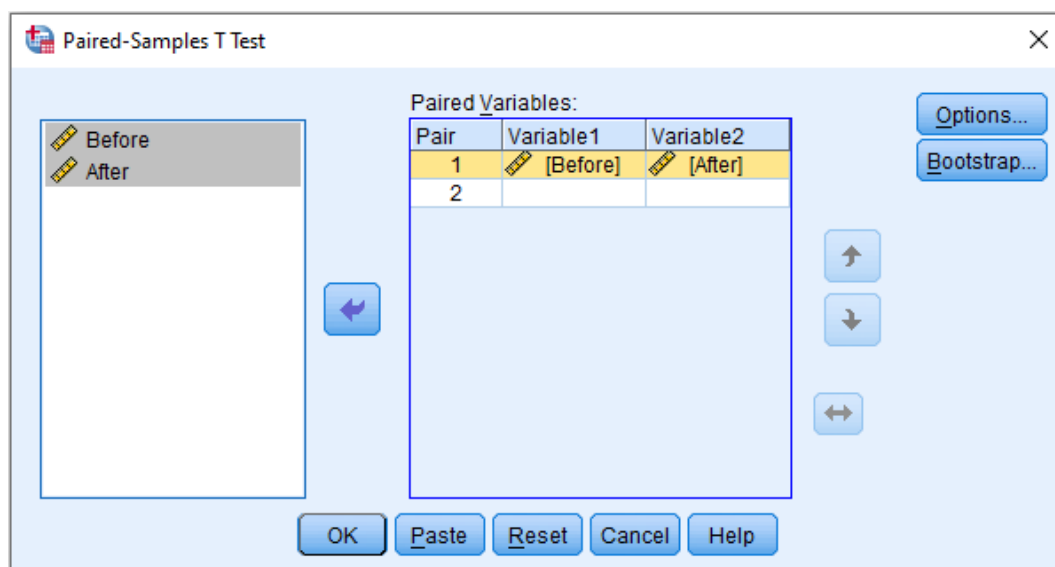


ภาพที่ 6.8 เมนูหน้าจอคำสั่ง Analyze ➤ Compare Means ➤ Paired-Samples T Test ...



ภาพที่ 6.9 หน้าจอคำสั่ง Paired-Samples T Test

(3) เลือกตัวแปรที่ต้องการวิเคราะห์ใส่ใน Paired Variables ในที่นี้คือ ตัวแปร **Before** และ **After** โดยให้กดคีย์บอร์ด **shift ↑** ค้างไว้ คลิกเลือกตัวแปร Before และ After เพื่อจับคู่ แล้วกดย้ายจากกล่อง ด้านซ้ายมือไปกล่อง Paired Variables ที่อยู่ด้านขวามือ จะได้ดังภาพที่ 6.10 จากนั้นกด OK เพื่อวิเคราะห์ ข้อมูล ทั้งนี้ นักวิจัยสามารถวิเคราะห์เปรียบเทียบค่าเฉลี่ยของข้อมูลได้หลาย ๆ คู่พร้อมกัน



ภาพที่ 6.10 การเลือกตัวแปรแบบจับคู่ใส่ในกล่อง Paired Variables

(4) ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 6.11

T-Test

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	Before	121.2500	8	12.39528	4.38239
	After	113.5000	8	13.32023	4.70941

Paired Samples Correlations

		N	Correlation	Sig.
Pair 1	Before & After	8	.908	.002

Paired Samples Test

	Paired Differences					t	df	Sig. (2-tailed)
	Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
				Lower	Upper			
Pair 1 Before - After	7.75000	5.57418	1.97077	3.08987	12.41013	3.932	7	.006

ภาพที่ 6.11 หน้าจอผลลัพธ์จากการวิเคราะห์ด้วยวิธี Paired-samples t-test

จากผลการวิเคราะห์ในภาพที่ 6.11 แสดงผลลัพธ์ได้ 3 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Paired Samples Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของปริมาณคอเลสเทอรอลก่อนและหลังใช้ยา จากข้อมูลดังกล่าวอธิบายได้ดังนี้

N	คือ จำนวนคนไข้ของกลุ่มก่อนใช้ยาและหลังใช้ยา (กลุ่มเดียวกัน จำนวนจึงเท่ากัน)
Mean	คือ ค่าเฉลี่ยปริมาณคอเลสเทอรอลก่อนใช้ยาและหลังใช้ยา
Std. Deviation	คือ ค่าส่วนเบี่ยงเบนมาตรฐานแสดงการกระจายของข้อมูลปริมาณคอเลสเทอรอลก่อนใช้ยาและหลังใช้ยา
Std. Error Mean	คือ ค่าความคลาดเคลื่อนมาตรฐานของข้อมูลปริมาณคอเลสเทอรอลก่อนใช้ยาและหลังใช้ยา

ส่วนที่ 2 ตาราง Paired Samples Correlations เป็นส่วนที่แสดงค่าสถิติสหสัมพันธ์สำหรับการทดสอบความสัมพันธ์ของข้อมูลทั้งสองกลุ่ม จากข้อมูลดังกล่าวอธิบายได้ดังนี้

Correlation	คือ ค่าสัมประสิทธิ์สหสัมพันธ์ของเพียร์สัน (r) แสดงถึงความสัมพันธ์ของข้อมูลทั้งสองกลุ่มที่นำมาทดสอบ จากตารางพบว่าค่า r มีค่าเท่ากับ 0.908 แสดงว่าปริมาณคอเลสเทอรอลก่อนและหลังใช้ยามีความสัมพันธ์กันสูงในทิศทางเดียวกัน
Sig.	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 เกี่ยวกับการทดสอบความสัมพันธ์ ดังนี้ H_0 : ปริมาณคอเลสเทอรอลก่อนและหลังใช้ยาไม่มีความสัมพันธ์กัน H_1 : ปริมาณคอเลสเทอรอลก่อนและหลังใช้ยามีความสัมพันธ์กัน

ส่วนที่ 3 ตาราง Paired Samples Test เป็นส่วนที่แสดงค่าสถิติสำหรับใช้ในการทดสอบค่าเฉลี่ยของข้อมูลทั้งสองกลุ่ม จากข้อมูลดังกล่าวอธิบายได้ ดังนี้

Mean	คือ ค่าเฉลี่ยของผลต่างระหว่างปริมาณคอเลสเทอรอลก่อนและหลังใช้ยา
Std. Deviation	คือ ค่าประมาณส่วนเบี่ยงเบนมาตรฐานของผลต่าง
Std. Error Mean	คือ ค่าความคลาดเคลื่อนมาตรฐานของผลต่าง
95% Confidence Interval of the Difference	คือ ค่าแสดงขอบเขตของช่วงความเชื่อมั่น 95% ของผลต่างค่าเฉลี่ย
t, df	คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง

Sig.(2-tailed)	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 เช่น กรณีทดสอบสมมติฐานสองหาง (Two-tailed test) สมมติฐานในการทดสอบมีดังนี้ H_0 : ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้ก่อนใช้ยา = ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้หลังใช้ยา H_1 : ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้ก่อนใช้ยา $\neq$ ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้หลังใช้ยา
----------------	--

6.3.3 การสรุปผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test

การสรุปผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test ใช้เกณฑ์เดียวกับวิธี One-sample t-test ในตารางที่ 6.2 จากการตั้งสมมติฐานทางสถิติแบบสองหาง (Two-tailed test) ที่นักวิจัยกำหนดไว้ ดังนี้

H_0 : ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้ก่อนใช้ยา = ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้หลังใช้ยา

H_1 : ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้ก่อนใช้ยา $\neq$ ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้หลังใช้ยา

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 จากผลลัพธ์ที่ได้ในภาพที่ 6.11 สามารถสรุปผลได้ดังนี้

เนื่องจากการทดสอบแบบสองหาง จะปฏิเสธสมมติฐาน H_0 เมื่อ ค่า Sig. (2 – tailed) $\leq \alpha$

การพิจารณาเพื่อสรุปผล จากตาราง Paired Samples Test ในภาพที่ 6.11 พบว่า ค่า Sig. (2-tailed) เท่ากับ 0.006 ซึ่งมีค่าน้อยกว่าค่า α ที่นักวิจัยกำหนดไว้ (ค่า $\alpha = 0.05$)

ดังนั้นจึงตัดสินใจได้ว่า ปฏิเสธสมมติฐานหลัก (Reject H_0)

สรุปผลได้ว่า ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้ก่อนใช้ยา $\neq$ ค่าเฉลี่ยปริมาณคอเลสเทอรอลของคนไข้หลังใช้ยาที่ระดับนัยสำคัญ 0.05

จากตัวอย่างการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test ข้างต้น เป็นการวิเคราะห์สมมติฐานแบบสองหาง (Two-tailed test) ถ้าในกรณีที่นักวิจัยต้องการวิเคราะห์สมมติฐานแบบหางเดียว (One-tailed test) เช่น ผู้วิจัยคาดว่าหลังใช้ยาไปแล้ว 1 เดือน คนไข้มีปริมาณคอเลสเทอรอลลดลง จะสามารถเขียนสมมติฐานทางสถิติแบบหางเดียว ได้ดังนี้

H_0 : ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยา $\leq$ ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้หลังใช้ยา

H_1 : ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยา $>$ ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้หลังใช้ยา

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 จากผลลัพธ์ที่ได้ในภาพที่ 6.11 สามารถสรุปผลได้ดังนี้

เนื่องจากการทดสอบทางเดียว-ทิศทางบวก จะปฏิเสธสมมติฐาน H_0 เมื่อ (1) ค่า t เป็นบวก และ (2) ค่า $\frac{\text{Sig.}(2\text{-tailed})}{2} \leq \alpha$ ถ้าข้อใดข้อหนึ่งไม่จริงจะไม่สามารถปฏิเสธสมมติฐาน H_0 ได้

การพิจารณาเพื่อสรุปผล จากตาราง Paired Samples Test ในภาพที่ 6.11 พบว่า ค่า Sig.(2-tailed) เท่ากับ 0.006 เมื่อนำค่าที่ได้มาหารสองจะได้ค่าเท่ากับ 0.003 ซึ่งมีค่าน้อยกว่าค่า α ที่นักวิจัยกำหนดไว้ (ค่า $\alpha = 0.05$) และมีค่า t เป็นบวก ($t = 3.932$)

ดังนั้นจึงตัดสินใจได้ว่า ปฏิเสธสมมติฐานหลัก (Reject H_0)

สรุปผลได้ว่า ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยา $>$ ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้หลังใช้ยาที่ระดับนัยสำคัญ 0.05

6.3.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test

จากภาพที่ 6.11 จะเห็นว่าในตาราง Paired Samples Test ในช่องสุดท้ายแสดง ค่า Sig. (2-tailed) เท่ากับ 0.006 ซึ่งมีค่าน้อยกว่าค่า α ที่นักวิจัยกำหนดไว้ (ค่า $\alpha = 0.05$) จึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้ก่อนใช้ยาไม่เท่ากับค่าเฉลี่ยปริมาณคอเลสเตอรอลของคนไข้หลังใช้ยาที่ระดับนัยสำคัญ 0.05 โดยก่อนใช้ยาคนไข้มีค่าเฉลี่ยปริมาณคอเลสเตอรอลเท่ากับ 121.25 ± 12.40 mg/dL และ หลังใช้ยามีค่าเฉลี่ยปริมาณคอเลสเตอรอลเท่ากับ 113.50 ± 13.32 mg/dL สามารถเขียนผลการวิเคราะห์ในรูปแบบตารางได้ดังแสดงในตารางที่ 6.6

ตารางที่ 6.6 ค่าเฉลี่ยปริมาณคอเลสเตอรอลก่อนใช้ยาและหลังใช้ยาของผู้ป่วย 8 ราย

ประเภทผู้ป่วย	ปริมาณคอเลสเตอรอล (mg/dL)
ก่อนใช้ยา	121.25 ± 12.40^a
หลังใช้ยา	113.50 ± 13.32^b

^{a-b} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

6.3.5 กรณีศึกษาการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test

นักพัฒนาผลิตภัณฑ์ของบริษัทแห่งหนึ่งต้องการพัฒนาสูตรเค้กไขมันต่ำ จึงเริ่มต้นการพัฒนาผลิตภัณฑ์โดยการคัดเลือกสูตรพื้นฐาน จากการรวบรวมข้อมูลทางอินเทอร์เน็ตได้สูตรเค้กมา 2 สูตร (สูตร A และสูตร B) จึงได้ทำการทดลองโดยให้ผู้ทดสอบจำนวน 30 คน ชิมเค้กทีละสูตรแล้วให้ประเมินความชอบด้วยสเกล 7-point hedonic scale (โดยคะแนน 1 = ไม่ชอบมากที่สุด และคะแนน 7 = ชอบมากที่สุด) ผลการทดสอบได้คะแนนความชอบโดยรวมแสดงดังตารางที่ 6.7

ตารางที่ 6.7 คะแนนความชอบเค้กสูตร A และสูตร B ของผู้ทดสอบจำนวน 30 คนประเมินความชอบด้วยสเกล 7-point hedonic scale (โดยคะแนน 1 = ไม่ชอบมากที่สุด และคะแนน 7 = ชอบมากที่สุด)

ผู้ทดสอบคนที่	เค้กสูตร A	เค้กสูตร B	ผู้ทดสอบคนที่	เค้กสูตร A	เค้กสูตร B
1	7	5	16	6	7
2	6	5	17	6	7
3	7	5	18	6	7
4	5	6	19	7	7
5	6	6	20	7	7
6	7	6	21	7	5
7	5	4	22	7	5
8	6	4	23	5	5
9	6	4	24	6	6
10	6	4	25	6	6
11	7	5	26	5	6
12	7	5	27	6	6
13	7	6	28	6	6
14	6	6	29	6	6
15	6	6	30	6	5

ขั้นตอนการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test

ขั้นที่ 1 การเขียนสมมติฐาน

จากตัวอย่าง การตั้งสมมติฐานเพื่อการทดสอบทำได้ 3 กรณีขึ้นอยู่กับวัตถุประสงค์ของผู้วิจัย ดังนี้

กรณีที่ 1 : ทดสอบสมมติฐานแบบสองหาง

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A = ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\neq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

กรณีที่ 2 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางมากกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\leq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $>$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

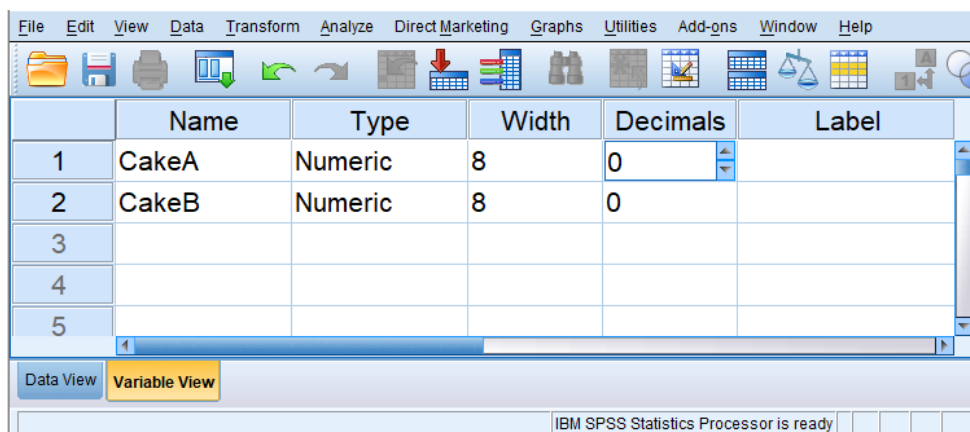
กรณีที่ 3 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางน้อยกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\geq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $<$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

ขั้นที่ 2 การวิเคราะห์ข้อมูลค่าเฉลี่ยด้วยวิธี Paired-samples t-test ในโปรแกรม SPSS

(1) เปิดโปรแกรม SPSS เพื่อป้อนข้อมูล ในกรณีนี้มีตัวแปร 2 ตัว คือ เค้กสูตร A และเค้กสูตร B ดังนั้นกำหนดชื่อตัวแปรเป็น **CakeA** และ **CakeB** ในหน้าต่าง Variable View ดังภาพที่ 6.12 จากนั้นป้อนข้อมูลคะแนนความชอบเค้กสูตร A และสูตร B ของผู้ทดสอบแต่ละคนโดยป้อนข้อมูลเป็นคู่ ๆ ในหน้าต่าง Data View ดังภาพที่ 6.13



ภาพที่ 6.12 กำหนดชื่อตัวแปร CakeA และ CakeB ในหน้าต่าง Variable View

The screenshot shows the Data View window in IBM SPSS. It contains a table with columns: CakeA, CakeB, and several empty columns labeled 'var'. There are 8 rows of data. The data for CakeA and CakeB is as follows:

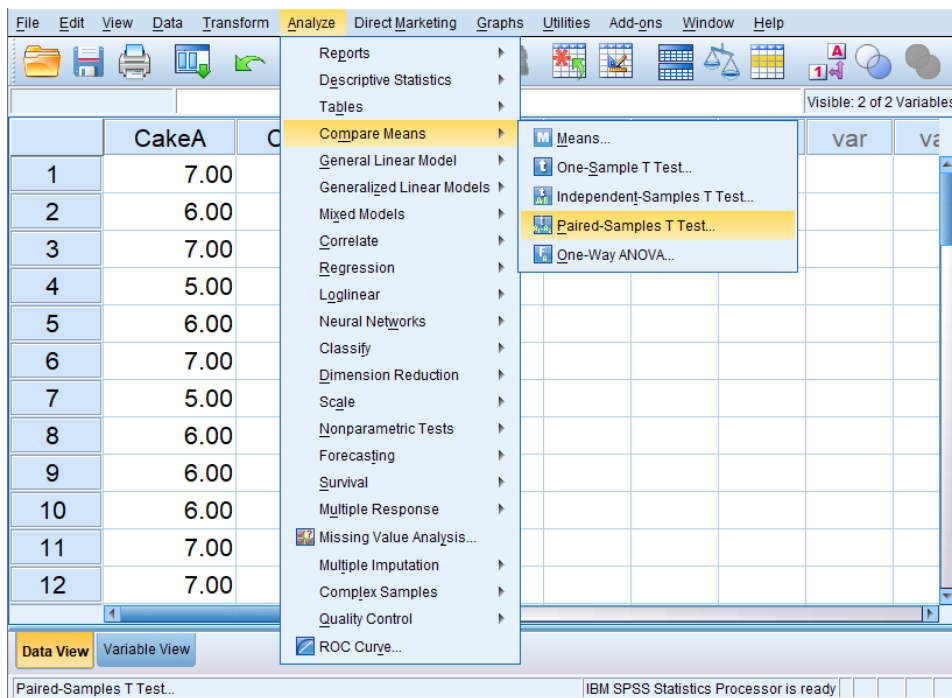
	CakeA	CakeB
1	7	5
2	6	5
3	7	5
4	5	6
5	6	6
6	7	6
7	5	4
8	6	4

The 'Data View' button is selected at the bottom.

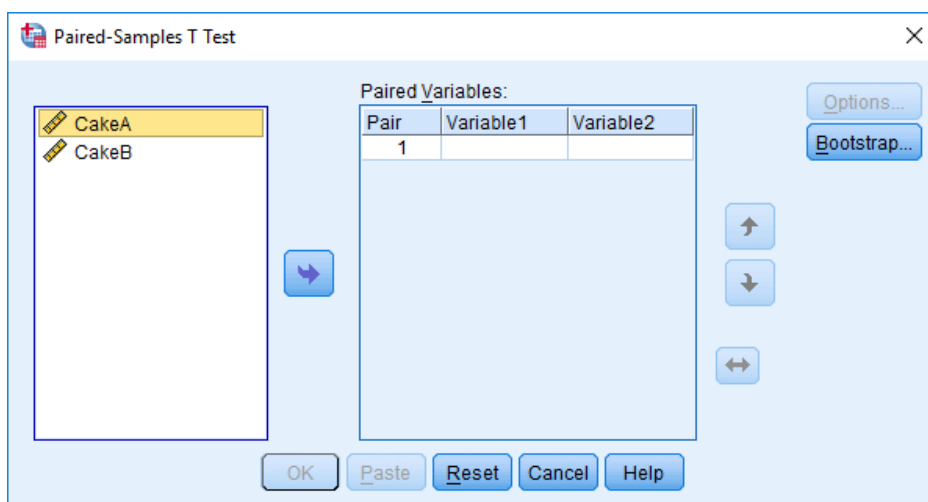
ภาพที่ 6.13 ป้อนข้อมูลคะแนนความชอบเป็นคู่ ๆ ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > Compare Means > Paired-Samples T Test...** ดังภาพที่ 6.14

จะปรากฏหน้าต่างดังภาพที่ 6.15



ภาพที่ 6.14 เมนูหน้าจอคำสั่ง Analyze > Compare Means > Paired-Samples T Test ...



ภาพที่ 6.15 หน้าจอคำสั่ง Paired-Samples T Test

(3) เลือกตัวแปร **CakeA** และ **CakeB** ที่ต้องการวิเคราะห์ใส่ใน Paired Variables โดยให้กดคีย์บอร์ด **shift ↑** ค้างไว้ คลิกเลือกตัวแปร **CakeA** และ **CakeB** เพื่อจับคู่ แล้วกดย้ายจากกล่องด้านซ้ายมือ ไปกล่อง Paired Variables ที่อยู่ด้านขวามือ จากนั้นกดคลิก OK จะได้ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 6.16

T-Test

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	CakeA	6.20	30	.664	.121
	CakeB	5.60	30	.932	.170

Paired Samples Correlations

		N	Correlation	Sig.
Pair 1	CakeA & CakeB	30	.078	.682

Paired Samples Test

		Paired Differences					t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair 1	CakeA - CakeB	.600	1.102	.201	.189	1.011	2.983	29	.006

ภาพที่ 6.16 หน้าจอผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test

ขั้นที่ 3 การสรุปผลการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Paired-samples t-test

จากผลลัพธ์ในภาพที่ 6.16 สามารถสรุปผลการวิเคราะห์ได้ 3 กรณีตามสมมติฐานในการทดสอบ ดังนี้

กรณีที่ 1 : ทดสอบสมมติฐานแบบสองหาง

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A = ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\neq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

จากตาราง Paired Samples Test ในช่องสุดท้ายแสดงค่า Sig. (2-tailed) เท่ากับ 0.006 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ จึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\neq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B ที่ระดับนัยสำคัญ 0.05

กรณีที่ 2 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางมากกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\leq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $>$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

จากตาราง Paired Samples Test ในช่องสุดท้ายแสดงค่า Sig. (2-tailed) เท่ากับ 0.006 เมื่อนำมาหารสองได้ค่าเท่ากับ 0.003 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ และค่า t มีค่าเป็นบวก ($t = 2.983$) ดังนั้น จึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $>$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B ที่ระดับนัยสำคัญ 0.05

กรณีศึกษาที่ 3 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางน้อยกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\geq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $<$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B

จากตาราง Paired Samples Test ในช่องสุดท้ายแสดงค่า Sig. (2-tailed) เท่ากับ 0.006 เมื่อนำมาหารสองได้ค่าเท่ากับ 0.003 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ และค่า t มีค่าเป็นบวก ($t = 2.983$) ดังนั้น จึงสรุปได้ว่า ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A $\geq$ ค่าเฉลี่ยคะแนนความชอบเค้กสูตร B ที่ระดับนัยสำคัญ 0.05

ขั้นที่ 4 การนำเสนอผลการวิเคราะห์ทางสถิติในรูปแบบตาราง

จากผลลัพธ์ที่แสดงในภาพที่ 6.16 จึงสรุปได้ว่า ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A ไม่เท่ากับค่าเฉลี่ยคะแนนความชอบเค้กสูตร B ที่ระดับนัยสำคัญ 0.05 โดยผู้ทดสอบให้คะแนนความชอบเค้กสูตร A เฉลี่ย เท่ากับ 6.2 ± 0.7 และ ให้คะแนนความชอบเค้กสูตร B เฉลี่ยเท่ากับ 5.6 ± 0.9 สามารถเขียนผลการวิเคราะห์ในรูปแบบตารางได้ดังแสดงในตารางที่ 6.8

ตารางที่ 6.8 ค่าเฉลี่ยคะแนนความชอบเค้กสูตร A และสูตร B ประเมินโดยผู้ทดสอบจำนวน 30 คน ด้วยสเกล 7-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 7 = ชอบมากที่สุด)

ตัวอย่างเค้ก	คะแนนความชอบโดยรวม
สูตร A	$6.2 \pm 0.7a$
สูตร B	$5.6 \pm 0.9b$

^{a-b}ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

จากการวิเคราะห์เปรียบเทียบค่าเฉลี่ยคะแนนความชอบเค้กสูตร A และสูตร B และพบว่าค่าเฉลี่ยคะแนนความชอบเค้กสูตร A ไม่เท่ากับค่าเฉลี่ยคะแนนความชอบเค้กสูตร B ที่ระดับนัยสำคัญ 0.05 ดังนั้น นักพัฒนาผลิตภัณฑ์ควรเลือกเค้กสูตร A เป็นสูตรพื้นฐานเพื่อนำไปพัฒนาเป็นเค้กสูตรใหม่ในขั้นตอนต่อไปของกระบวนการพัฒนาผลิตภัณฑ์ เนื่องจากเมื่อพิจารณาค่าเฉลี่ยพบว่าสูตร A มีค่าเฉลี่ยคะแนนความชอบมากกว่าสูตร B

6.4 การทดสอบสมมติฐานทางสถิติเกี่ยวกับค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่างที่เป็นอิสระกัน

วิธี Independent-samples t-test หรือ Two samples t-test เป็นการเปรียบเทียบค่าเฉลี่ยของ 2 กลุ่มตัวอย่างซึ่งเป็นอิสระจากกันหรือจากประชากรต่างกลุ่มกันว่ามีความแตกต่างกันหรือไม่ วิธี Independent-samples t-test นี้จำนวนข้อมูลของแต่ละกลุ่มไม่จำเป็นต้องเท่ากัน ตัวอย่างการทดสอบค่าเฉลี่ย 2 กลุ่มตัวอย่างที่เป็นอิสระกัน เช่น การทดสอบความชอบผลิตภัณฑ์ชาลาเปาสีลาวาของนิสิตชายและนิสิตหญิง การทดสอบความพึงพอใจของผู้บริโภคที่อยู่ในวัยเรียนและวัยทำงาน

6.4.1 การคำนวณค่าสถิติ t ด้วยวิธี Independent-samples t-test

วิธี Independent-samples t-test นี้จัดอยู่ในประเภทของการวิเคราะห์ข้อมูลแบบ 2 ตัวแปร (Bivariate data analysis) เนื่องจากในการวิเคราะห์ต้องใช้ตัวแปร 2 ตัว คือ (1) **ตัวแปรต้น** จะต้องเป็นตัวแปรที่ใช้สเกลที่สามารถจำแนกกลุ่มได้ 2 กลุ่ม ได้แก่ สเกลนามกำหนด เช่น เพศ (เพศชาย และ เพศหญิง) อาชีพ (อาจารย์ และ นิสิต) หรือจะเป็นสเกลแบบช่วงและสเกลอัตราส่วนที่ได้ทำการจัดกลุ่มแล้ว เช่น อุณหภูมิในการทอด (100 และ 120 องศาเซลเซียส) และ (2) **ตัวแปรตาม** จะต้องเป็นตัวแปรที่แสดงคุณลักษณะของกลุ่มตัวอย่างและสามารถนำมาใช้คำนวณได้ นั่นคือ สเกลแบบช่วง และสเกลอัตราส่วน ทั้งนี้ข้อมูลที่สามารถใช้วิธี Independent-samples t-test วิเคราะห์เปรียบเทียบค่าเฉลี่ยได้จะต้องเป็นข้อมูลที่มีการแจกแจงแบบปกติหรือใกล้เคียงปกติ ไม่จำเป็นต้องทราบการกระจายของประชากรทั้ง 2 กลุ่ม (ไม่จำเป็นต้องทราบ σ_1 , σ_2) แต่ต้องนำข้อมูลที่ได้จากตัวอย่างมาทดสอบก่อนเพื่อพิจารณาว่าต้องใช้สูตรใดในการคำนวณค่าสถิติ t

การคำนวณค่าสถิติ t มีสูตรในการคำนวณที่แตกต่างกัน ขึ้นอยู่กับเงื่อนไข ดังนี้

(1) เมื่อการกระจายของประชากรทั้ง 2 กลุ่มแตกต่างกัน ($\sigma_1 \neq \sigma_2$)

$$t = \frac{(\bar{x}_1 + \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}, \quad Df = \frac{(S_1^2/n_1 + S_2^2/n_2)}{(\left(\frac{S_1^2}{n_1}\right)^2 / (n_1 - 1)) + (\left(\frac{S_2^2}{n_2}\right)^2 / (n_2 - 1))}$$

(2) เมื่อการกระจายของประชากรทั้ง 2 กลุ่มไม่แตกต่างกัน ($\sigma_1 = \sigma_2$)

$$t = \frac{(\bar{x}_1 + \bar{x}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{S_p^2}{n_1} + \frac{S_p^2}{n_2}}}, \quad S_p^2 = \frac{(n_1 - 1)S_1^2 + (n_2 - 1)S_2^2}{(n_1 + n_2 - 2)}$$

S_p^2 (Pooled Variance Estimate) เป็นค่าประมาณของความแปรปรวนร่วมทั้ง 2 กลุ่ม

6.4.2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

ตัวอย่างที่ 1 นักวิจัยต้องการทราบว่าผู้ทดสอบชิมที่เป็นนิสิตเรียนสายวิทยาศาสตร์ (30 คน) กับ นิสิตที่เรียนสายสังคมศาสตร์ (26 คน) ชอบผลิตภัณฑ์น้ำสลัดผักทองที่ถูกพัฒนาขึ้นมาใหม่แตกต่างกันหรือไม่ จึงให้ผู้ทดสอบทั้งสองกลุ่มทำการทดสอบตัวอย่างน้ำสลัดผักทองด้วยสเกล 9-point hedonic scale (โดยคะแนน 1 = ไม่ชอบมากที่สุด และ คะแนน 9 = ชอบมากที่สุด) โดยมีผลคะแนนความชอบแสดงดังตารางที่ 6.9

ตารางที่ 6.9 คะแนนความชอบผลิตภัณฑ์น้ำสลัดผักทองของนิสิตที่เรียนสายวิทยาศาสตร์และสายสังคมศาสตร์ ประเมินด้วยสเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด)

ผู้ทดสอบ คนที่	คะแนนความชอบ	
	นิสิตสายวิทย์	นิสิตสายสังคม
1	7	6
2	6	6
3	7	6
4	6	6
5	6	6
6	6	7
7	6	7
8	7	7
9	7	8
10	7	8
11	7	9
12	7	9
13	7	6
14	6	6
15	6	6

ผู้ทดสอบ คนที่	คะแนนความชอบ	
	นิสิตสายวิทย์	นิสิตสายสังคม
16	6	7
17	6	7
18	6	7
19	7	7
20	7	7
21	7	8
22	7	8
23	6	7
24	6	6
25	6	6
26	5	7
27	6	
28	6	
29	6	
30	6	

จากข้อมูลข้างต้น พิจารณารายละเอียดของข้อมูลได้ ดังนี้

- (1) ตัวแปรต้น คือ กลุ่มนิสิต (มี 2 กลุ่ม ได้แก่ นิสิตที่เรียนสายวิทยาศาสตร์ และนิสิตที่เรียนสายสังคมศาสตร์ ข้อสังเกต นิสิตทั้ง 2 กลุ่มเป็นกลุ่มตัวอย่างที่มีความเป็นอิสระกัน)
- (2) ตัวแปรตาม คือ คะแนนความชอบ

การตั้งสมมติฐานทางสถิติเพื่อการทดสอบ ทำได้ 3 กรณีขึ้นอยู่กับวัตถุประสงค์ผู้วิจัย ดังนี้

กรณีที่ 1 : ทดสอบสมมติฐานแบบสองหาง

H_0 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ = ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

H_1 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $\neq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

กรณีที่ 2 : ทดสอบสมมติฐานแบบทางเดียว – ทิศทางมากกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $\leq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

H_1 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $>$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

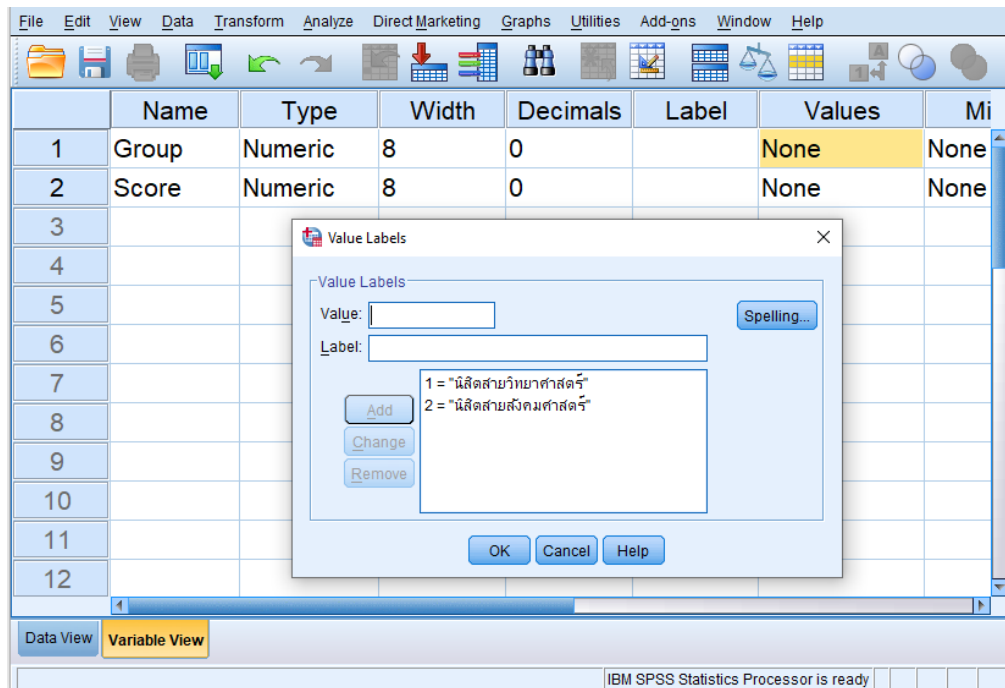
กรณีที่ 3 : ทดสอบสมมติฐานแบบทางเดียว – ทิศทางน้อยกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $\geq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

H_1 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $<$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

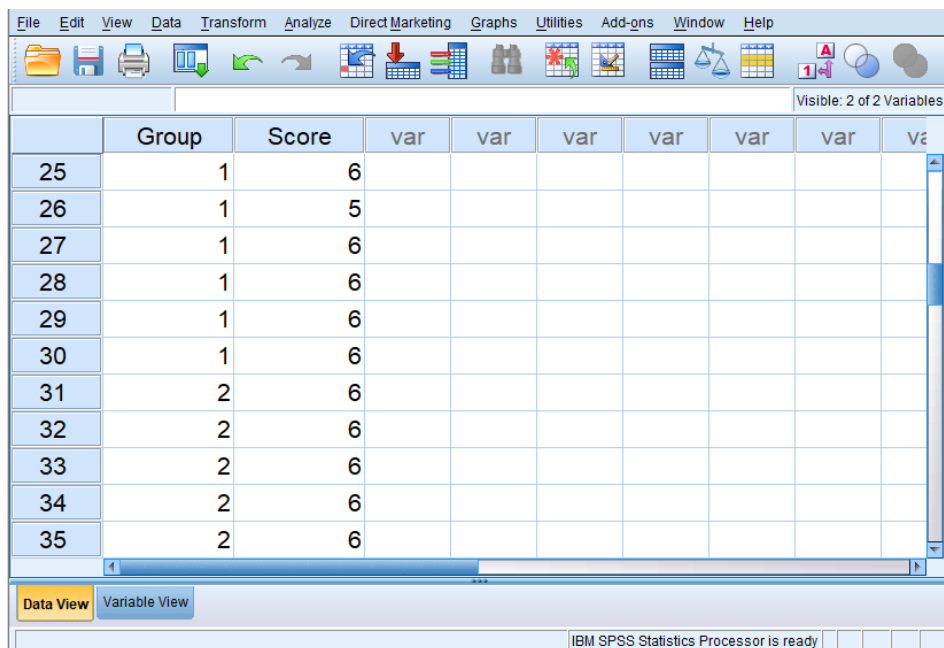
ขั้นตอนการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test โดยโปรแกรม SPSS

(1) เปิดโปรแกรม SPSS เพื่อป้อนข้อมูล ในกรณีนี้จะมีตัวแปร 2 ตัว คือ ตัวแปรกลุ่มนิสิต และ ตัวแปรคะแนนความชอบ ดังนั้นกำหนดชื่อตัวแปรเป็น **Group** และ **Score** ในหน้าต่าง Variable View กำหนดทศนิยมทั้งสองตัวแปรให้เท่ากับ 0 และกำหนดค่า Values ให้กับตัวแปร Group โดยกำหนดให้ 1 คือ นิสิตสายวิทยาศาสตร์ และ 2 คือ นิสิตสายสังคมศาสตร์ ดังภาพที่ 6.17 (วิธีการกำหนดค่า Values ให้ตัวแปรสามารถอ่านเพิ่มเติมได้ในบทที่ 3)



ภาพที่ 6.17 การกำหนดชื่อตัวแปรและค่าตัวแปร Group ในหน้าต่าง Variable View

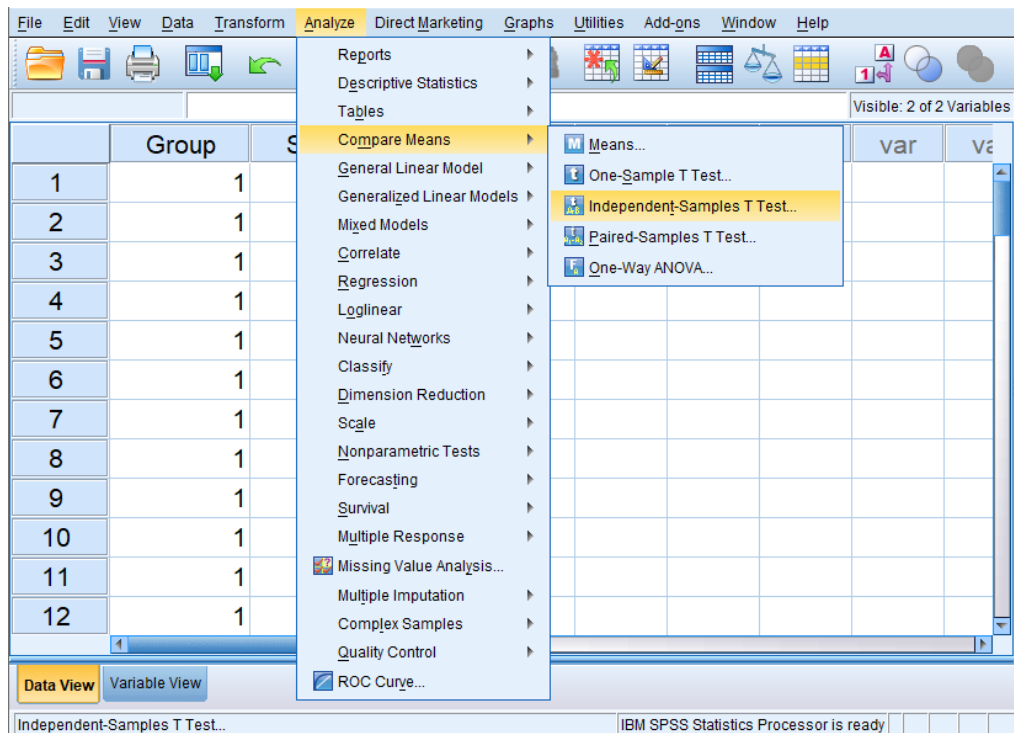
(2) จากนั้นป้อนข้อมูลกลุ่ม (ตัวแปร Group) และ คะแนนความชอบ (ตัวแปร Score) ลงในหน้าต่าง Data View ดังภาพที่ 6.18 แสดงให้เห็นการป้อนข้อมูลจะป้อนชุดข้อมูลของกลุ่มนิสิตสายวิทยาศาสตร์ (Group 1) ให้เสร็จก่อน (Case ที่ 1-30) และต่อด้วยชุดข้อมูลกลุ่มนิสิตสายสังคมศาสตร์ (Group 2) (Case ที่ 31-56)



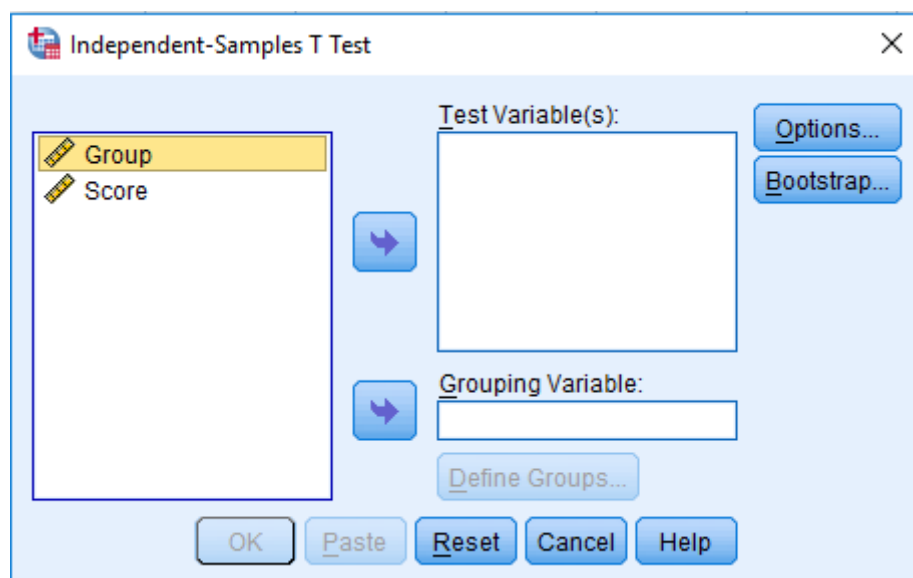
ภาพที่ 6.18 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > Compare Means > Independent-Samples T Test...** ดัง

ภาพที่ 6.19 จะปรากฏหน้าต่างดังภาพที่ 6.20

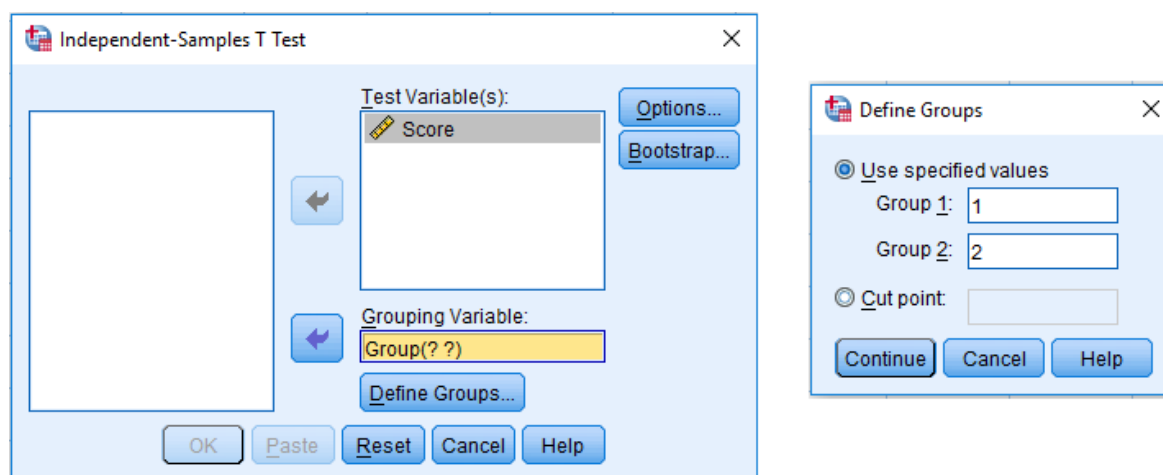


ภาพที่ 6.19 เมนูหน้าต่างคำสั่ง Analyze > Compare Means > Independent-Samples T Test...



ภาพที่ 6.20 หน้าจอคำสั่ง Independent-Samples T Test

(4) เลือกตัวแปร **Score** ใส่ในกล่อง Test Variable(s) และตัวแปร **Group** ใส่ในกล่อง Grouping Variable จากนั้นให้กด **Define Groups...** เพื่อกำหนดเลขกลุ่ม ในที่นี้คือ 1 และ 2 (เนื่องจากกำหนดให้ 1 คือ นิสิตสายวิทยาศาสตร์ และ 2 คือ นิสิตสายสังคมศาสตร์) จะได้ดังภาพที่ 6.21



ภาพที่ 6.21 เลือกตัวแปร Score ใส่ในกล่อง Test Variable(s) และตัวแปร Group ใส่ในกล่อง Grouping Variable แล้วกด Define Groups เพื่อกำหนดเลขกลุ่ม (ในที่นี้คือ 1 และ 2)

(5) ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 6.22

T-Test

Group Statistics				
Group	N	Mean	Std. Deviation	Std. Error Mean
Score นิสิตสายวิทยาศาสตร์	30	6.37	.556	.102
นิสิตสายสังคมศาสตร์	26	6.92	.935	.183

Independent Samples Test									
		Levene's Test for Equality of Variances		t-test for Equality of Means					
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference
Score	Equal variances assumed	3.111	.083	-2.749	54	.008	-.556	.202	Lower: -.962 Upper: -.151
	Equal variances not assumed			-2.655	39.483	.011	-.556	.210	Lower: -.980 Upper: -.133

ภาพที่ 6.22 หน้าจอผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

จากผลการวิเคราะห์ในภาพที่ 6.22 แสดงผลลัพธ์ได้ 2 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Group Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบ ในกรณีนี้คือ ตัวแปร Score จากข้อมูลดังกล่าวอธิบายได้ดังนี้

Group	คือ กลุ่มตัวอย่างตามที่ได้กำหนดค่า Values ไว้ (หากนักวิจัยไม่ได้กำหนดค่าที่แท้จริงให้กับรหัส ตรงส่วนนี้จะปรากฏเป็นตัวเลขรหัสของกลุ่ม นั่นคือเลข 1 และ 2)
N	คือ ค่าแสดงจำนวนของแต่ละกลุ่ม
Mean	คือ ค่าเฉลี่ยคะแนนความชอบของนิสิตแต่ละกลุ่ม
Std. Deviation	คือ ค่าส่วนเบี่ยงเบนมาตรฐานแสดงการกระจายของคะแนนความชอบนิสิตแต่ละกลุ่ม
Std. Error Mean	คือ ค่าความคลาดเคลื่อนมาตรฐานของข้อมูลคะแนนความชอบนิสิตแต่ละกลุ่ม

ส่วนที่ 2 ตาราง Independent Samples Test เป็นส่วนที่แสดงค่าสถิติสำหรับทดสอบค่าตัวแปร Score ซึ่งวัตถุประสงค์หลัก คือ เพื่อทดสอบค่าเฉลี่ยของข้อมูล 2 กลุ่ม แต่การทดสอบมีเงื่อนไขที่ต้องพิจารณา คือ การกระจายของข้อมูลทั้ง 2 กลุ่มแตกต่างกันหรือไม่แตกต่างกัน

2.1 ตาราง Levene's Test for Equality of Variances ใช้พิจารณาการกระจายของข้อมูล

F	คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในกรณีนี้การทดสอบสมมติฐานเป็นแบบสองหาง (Two-tailed test) เขียนได้ดังนี้ $H_0 : \text{ข้อมูลทั้ง 2 กลุ่มมีการกระจายตัวไม่แตกต่างกัน } (\sigma_1 = \sigma_2)$ $H_1 : \text{ข้อมูลทั้ง 2 กลุ่มมีการกระจายตัวแตกต่างกัน } (\sigma_1 \neq \sigma_2)$

2.2 ตาราง t-test for Equality of Means ใช้พิจารณาค่าเฉลี่ยของ 2 กลุ่มตัวอย่าง

t, df	คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.(2-tailed)	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 กรณีนี้ การทดสอบสมมติฐานเป็นแบบสองหาง (Two-tailed test) สมมติฐานเขียนได้ ดังนี้ $H_0 : \text{ค่าเฉลี่ยข้อมูลกลุ่มที่ 1} = \text{ค่าเฉลี่ยข้อมูลกลุ่มที่ 2}$ $H_1 : \text{ค่าเฉลี่ยข้อมูลกลุ่มที่ 1} \neq \text{ค่าเฉลี่ยข้อมูลกลุ่มที่ 2}$
Mean Difference	คือ ค่าผลต่างระหว่างค่าเฉลี่ยทั้ง 2 กลุ่ม
Std. Error Difference	คือ ค่าความคลาดเคลื่อนมาตรฐานของค่าผลต่าง

95% Confidence Interval of the Difference คือ ค่าแสดงขอบเขตของช่วงความเชื่อมั่น 95% ของผลต่างค่าเฉลี่ย

6.4.3 การสรุปผลลัพท์การวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

เพื่อสรุปผลการทดสอบค่าเฉลี่ยของข้อมูล 2 กลุ่ม นักวิจัยต้องพิจารณาค่าทางสถิติที่ใช้ทดสอบการกระจายของข้อมูลทั้ง 2 กลุ่มว่ามีความแตกต่างหรือไม่แตกต่างกัน เนื่องจากค่าสถิติ t ใช้สูตรคำนวณต่างกัน ขั้นตอนในการสรุปผลลัพท์การวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test มี 2 ขั้นตอน ดังนี้

ขั้นที่ 1 พิจารณาการกระจายของข้อมูลทั้ง 2 กลุ่ม แตกต่างหรือไม่

เขียนสมมติฐานในการทดสอบได้ ดังนี้

H_0 : ข้อมูลทั้งสองกลุ่มมีการกระจายไม่แตกต่างกัน ($\sigma_1 = \sigma_2$)

H_1 : ข้อมูลทั้งสองกลุ่มมีการกระจายแตกต่างกัน ($\sigma_1 \neq \sigma_2$)

พิจารณาค่า Sig. จากตาราง Levene's Test for Equality of Variance พบว่า ค่า Sig. เท่ากับ 0.083 ซึ่งมีค่ามากกว่าค่า α ที่นักวิจัยกำหนดไว้ (ค่า $\alpha = 0.05$) ดังนั้นจึงสรุปได้ว่า ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้งสองกลุ่มมีการกระจายไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05 หรือ “Equal variances assumed”

ขั้นที่ 2 พิจารณาค่าเฉลี่ยของทั้ง 2 กลุ่ม แตกต่างหรือไม่

ให้พิจารณาค่า Sig. (2-tailed) จากตาราง t-test for Equality of Means และสมมติฐานทางสถิติที่ตั้งไว้ ทั้งนี้เมื่อพิจารณาจะพบว่าข้อมูลในส่วน of ตาราง t-test for Equality of Means มี 2 บรรทัด คือ แถวบน และแถวล่าง ในการจะเลือกใช้ค่าสถิติแถวใดขึ้นอยู่กับผลการวิเคราะห์ที่ได้จากขั้นตอนที่ 1 ดังนี้

(1) หากพบว่าข้อมูลทั้งสองกลุ่มมีการกระจายไม่แตกต่างกัน (Equal variances assumed) ให้พิจารณาข้อมูลสถิติในส่วน of ตาราง t-test for Equality of Means แถวบน

(2) หากพบว่าข้อมูลทั้งสองกลุ่มมีการกระจายแตกต่างกัน (Equal variances not assumed) ให้พิจารณาข้อมูลสถิติในส่วน of ตาราง t-test for Equality of Means แถวล่าง

จากขั้นตอนที่ 1 พบว่า ข้อมูลทั้งสองกลุ่มมีการกระจายไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05 หรือ “Equal variances assumed” ดังนั้นในขั้นตอนที่ 2 จะพิจารณาข้อมูลสถิติในส่วน of ตาราง t-test for Equality of Means แถวบน นั่นคือ ค่า Sig. (2-tailed) เท่ากับ 0.008

จากผลลัพธ์ในภาพที่ 6.22 สามารถสรุปผลการวิเคราะห์ได้ 3 กรณีตามสมมติฐานในการทดสอบ โดยใช้เกณฑ์ในการปฏิเสธสมมติฐานหลัก (H_0) ดังตารางที่ 6.2 ดังนี้

กรณีที่ 1 : ทดสอบสมมติฐานแบบสองหาง

H_0 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ = ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

H_1 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $\neq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

จากตาราง t-test for Equality of Means แถวบน แสดง ค่า Sig. (2-tailed) เท่ากับ 0.008 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ จึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทยาศาสตร์ $\neq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคมศาสตร์ ที่ระดับนัยสำคัญ 0.05

กรณีที่ 2 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางมากกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $\leq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

H_1 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $>$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

จากตาราง t-test for Equality of Means แถวบน แสดง ค่า Sig. (2-tailed) เท่ากับ 0.008 เมื่อนำมาหารสองได้ค่าเท่ากับ 0.004 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ และค่า t มีค่าเป็นลบ ($t = -2.749$) ดังนั้นจึงสรุปได้ว่า ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทยาศาสตร์ $\leq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคมศาสตร์ ที่ระดับนัยสำคัญ 0.05

กรณีที่ 3 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางน้อยกว่า

H_0 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $\geq$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

H_1 : ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทย์ $<$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคม

จากตาราง t-test for Equality of Means แถวบน แสดง ค่า Sig. (2-tailed) เท่ากับ 0.008 เมื่อนำมาหารสองได้ค่าเท่ากับ 0.004 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ และค่า t มีค่าเป็นลบ ($t = -2.749$) ดังนั้นจึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทยาศาสตร์ $<$ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคมศาสตร์ ที่ระดับนัยสำคัญ 0.05

6.4.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

จากภาพที่ 6.22 สรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทยาศาสตร์ไม่เท่ากับค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคมศาสตร์ที่ระดับนัยสำคัญ 0.05 โดยค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทยาศาสตร์ เท่ากับ 6.4 ± 0.6 และ ค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคมศาสตร์ เท่ากับ 6.9 ± 0.9 สามารถเขียนผลการวิเคราะห์ในรูปแบบตารางได้ดังแสดงในตารางที่ 6.10

ตารางที่ 6.10 ค่าเฉลี่ยคะแนนความชอบน้ำสลัดผักทองของนิสิตที่เรียนสายวิทยาศาสตร์และสายสังคมศาสตร์ ประเมินด้วยสเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด)

กลุ่มนิสิต	คะแนนความชอบโดยรวม
นิสิตสายวิทยาศาสตร์ (n=30)	$6.4 \pm 0.6b$
นิสิตสายสังคมศาสตร์ (n=26)	$6.9 \pm 0.9a$

^{a-b}ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

จากการวิเคราะห์เปรียบเทียบค่าเฉลี่ยคะแนนความชอบน้ำสลัดผักทองของนิสิตที่เรียนสายวิทยาศาสตร์และสายสังคมศาสตร์ พบว่า ค่าเฉลี่ยคะแนนความชอบของนิสิตสายวิทยาศาสตร์ไม่เท่ากับค่าเฉลี่ยคะแนนความชอบของนิสิตสายสังคมศาสตร์ที่ระดับนัยสำคัญ 0.05 โดยที่นิสิตสายสังคมศาสตร์ให้คะแนนความชอบน้ำสลัดผักทองมากกว่านิสิตสายวิทยาศาสตร์

6.4.5 กรณีศึกษาการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

นักวิจัยต้องการเปรียบเทียบวิธีการสกัดสารชนิดหนึ่งโดยใช้วิธีมาตรฐาน (วิธี A) และวิธีแบบรวดเร็ว (วิธี B) โดยวิธี A เป็นวิธีการมาตรฐานที่มีความถูกต้องแต่ใช้เวลาในการสกัดสารนาน ส่วนวิธี B มีความรวดเร็วเนื่องจากในการสกัดสารใช้เวลาที่สั้นกว่า นักวิจัยต้องการนำวิธี B มาใช้แทนวิธี A จึงต้องการทราบว่าทั้งสองวิธีสกัดสารดังกล่าวได้ปริมาณสารสกัด (%Yield) แตกต่างกันหรือไม่ นักวิจัยนำตัวอย่างรำข้าวมาสกัดด้วยวิธี A (ทำการสกัด 4 ซ้ำ) และวิธี B (ทำการสกัด 7 ซ้ำ) ได้ผลการวิเคราะห์ดังแสดงในตารางที่ 6.11

ตารางที่ 6.11 ปริมาณสารที่สกัดได้ (%Yield) ในตัวอย่างรำข้าวโดยใช้การสกัดวิธี A และวิธี B

วิธีการสกัดสาร	ปริมาณสารที่สกัดได้ (%)						
วิธี A	26.45	26.66	26.37	26.44			
วิธี B	25.95	26.35	25.78	26.11	26.05	25.87	25.98

จากตารางที่ 6.8 ซึ่งแสดงข้อมูลปริมาณสารที่สกัดได้ (%) ในตัวอย่างรำข้าวที่สกัดโดยวิธี A และวิธี B เมื่อสังเกตพบว่าจำนวนข้อมูล (n) ไม่เท่ากัน โดยวิธี A มีจำนวน 4 ค่า และวิธี B มีจำนวน 7 ค่า เนื่องจากระยะเวลาที่ใช้ในการวิเคราะห์สารไม่เท่ากันทำให้ข้อมูลที่ได้มีจำนวนต่างกัน

ขั้นที่ 1 การตั้งสมมติฐานทางสถิติเพื่อการทดสอบ ทำได้ 3 กรณีขึ้นอยู่กับวัตถุประสงค์ผู้วิจัย ดังนี้

กรณีที่ 1 : ทดสอบสมมติฐานแบบสองหาง

H_0	:	ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A = ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B
H_1	:	ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\neq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

กรณีที่ 2 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางมากกว่า

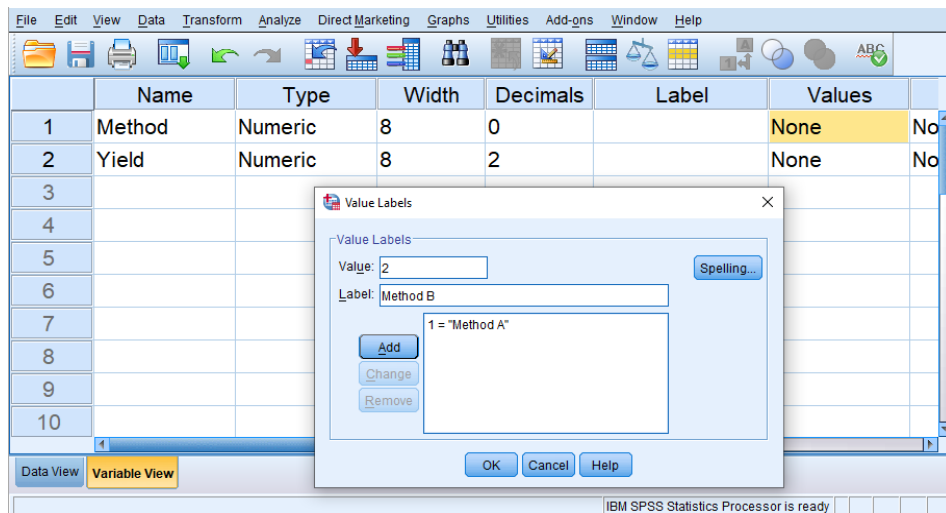
H_0	:	ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\leq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B
H_1	:	ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $>$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

กรณีที่ 3 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางน้อยกว่า

H_0	:	ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\geq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B
H_1	:	ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $<$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

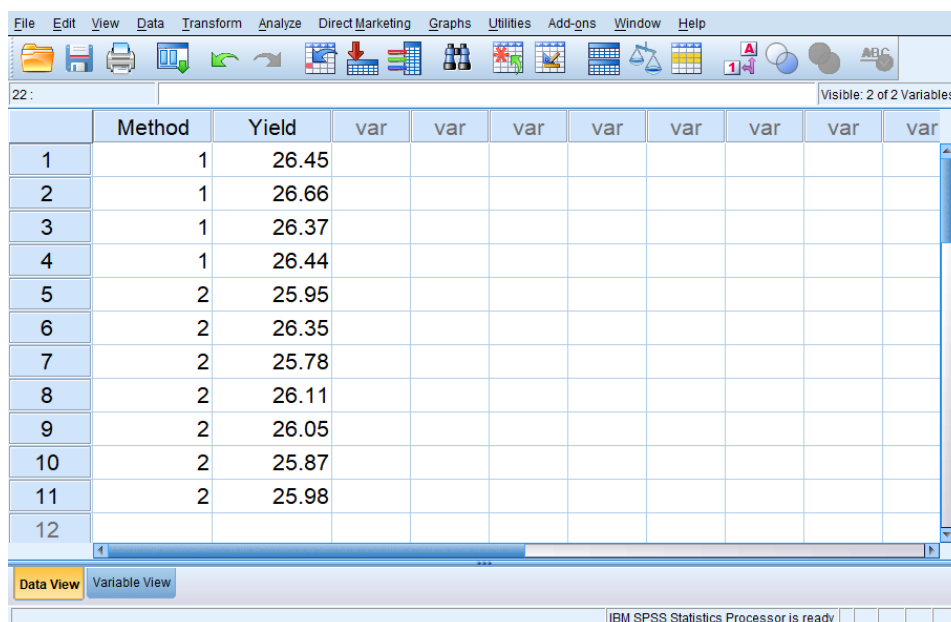
ขั้นที่ 2 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ Independent-samples t-test

(1) เปิดโปรแกรม SPSS แล้วป้อนข้อมูลผลการสกัดสารโดยวิธี A และวิธี B ในกรณีนี้จะมีตัวแปร 2 ตัว คือ วิธีการสกัด และ ปริมาณสารสกัด ดังนั้นกำหนดชื่อเป็น ตัวแปร **Method** และตัวแปร **Yield** ในหน้าต่าง Variable View และทำการกำหนดทศนิยมตัวแปร Yield ให้เท่ากับ 2 และตัวแปร Method ให้เท่ากับ 0 จากนั้นกำหนดค่า Values ให้กับตัวแปร Method โดยกำหนดให้ 1 คือ Method A และ 2 คือ Method B ดังภาพที่ 6.23



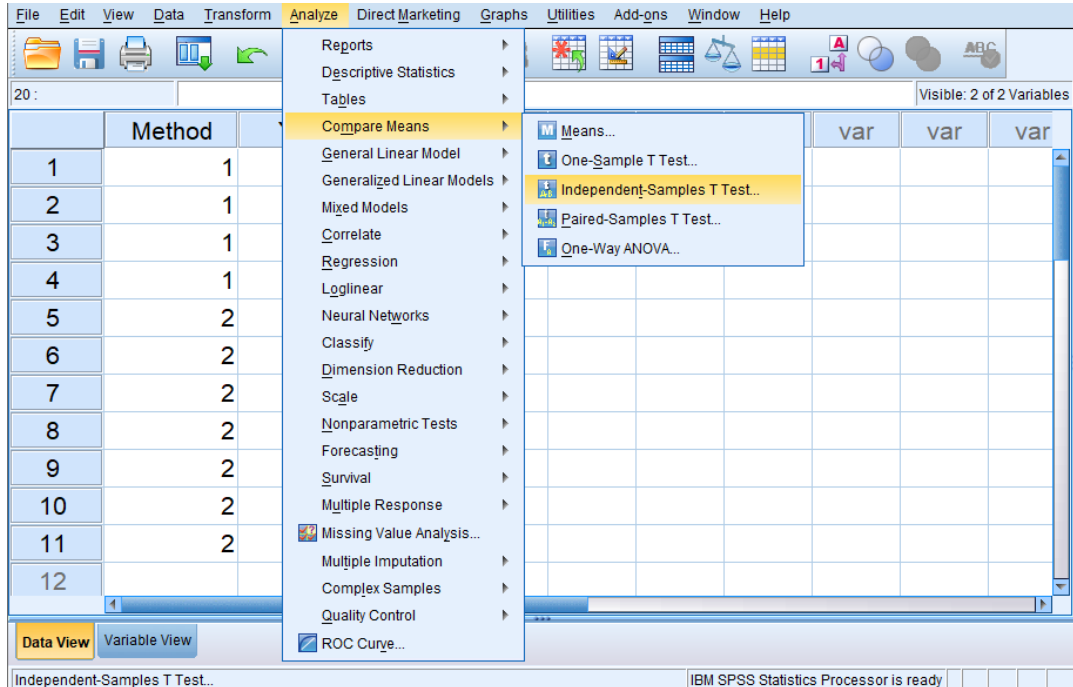
ภาพที่ 6.23 กำหนดชื่อ ตัวแปร Method และตัวแปร Yield ในหน้าต่าง Variable View

(2) ป้อนข้อมูลวิธีการสกัด และปริมาณสารสกัด ลงในหน้าต่าง Data View ดังภาพที่ 6.24

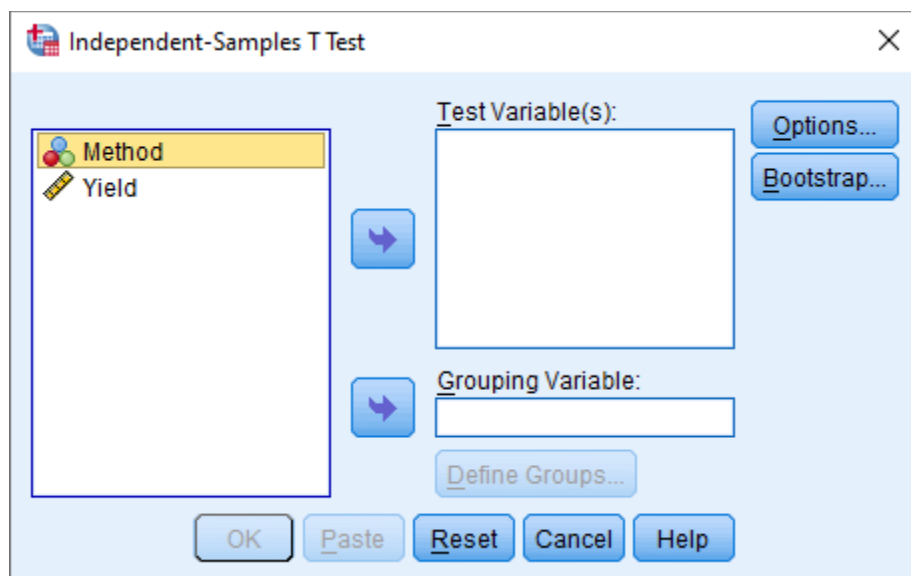


ภาพที่ 6.24 ป้อนข้อมูลวิธีการสกัด และปริมาณสารสกัด ลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > Compare Means > Independent-Samples T Test...** ดังภาพที่ 6.25 จะปรากฏหน้าต่างดังภาพที่ 6.26

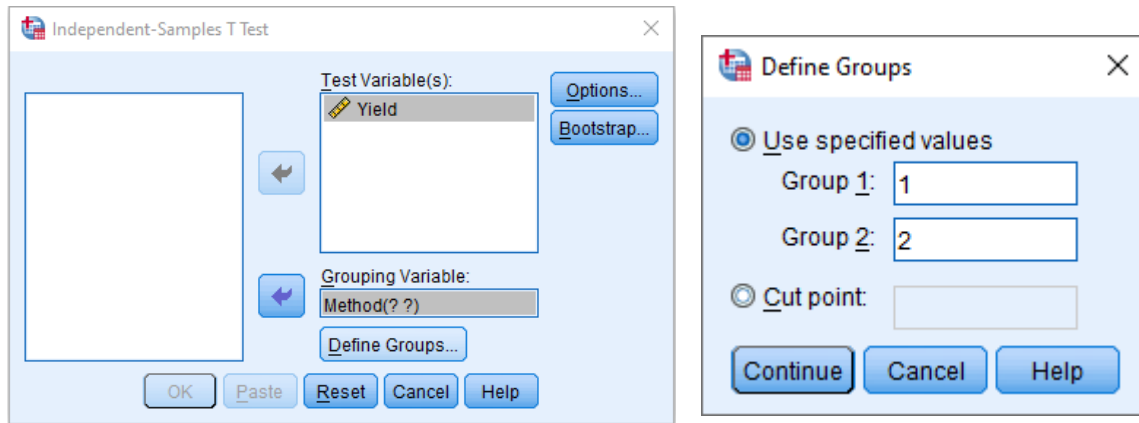


ภาพที่ 6.25 เมนูหน้าต่างคำสั่ง Analyze>Compare Means>Independent-Samples T Test...



ภาพที่ 6.26 หน้าจอคำสั่ง Independent-Samples T Test

(4) เลือกตัวแปร **Yield** ใส่ในกล่อง Test Variable(s) และตัวแปร **Method** ใส่ในกล่อง Grouping Variable โดยให้กด **Define Groups...** เพื่อกำหนดเลขกลุ่ม ในที่นี้คือ 1 และ 2 (เนื่องจากกำหนดให้ 1 คือ Method A และ 2 คือ Method B) จะได้ดังภาพที่ 6.27



ภาพที่ 6.27 เลือกตัวแปร Yield ใส่ในกล่อง Test Variable(s) และตัวแปร Method ใส่ในกล่อง Grouping Variable จากนั้นกด Define Groups เพื่อกำหนดเลขกลุ่ม (ในที่นี้คือ 1 และ 2)

(5) ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 6.28

T-Test

Group Statistics					
Method	N	Mean	Std. Deviation	Std. Error Mean	
Yield Method A	4	26.4800	.12517	.06258	
Method B	7	26.0129	.18446	.06972	

Independent Samples Test									
		Levene's Test for Equality of Variances		t-test for Equality of Means					
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference Lower Upper
Yield	Equal variances assumed	.499	.498	4.462	9	.002	.46714	.10470	.23029 .70400
	Equal variances not assumed			4.986	8.512	.001	.46714	.09369	.25334 .68094

ภาพที่ 6.28 หน้าจอผลลัพธ์จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

ขั้นที่ 3 การสรุปผลการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

จากผลการวิเคราะห์ในภาพที่ 6.28 มีขั้นตอนในการสรุปผล 2 ขั้นตอน ดังนี้

(1) ให้ตรวจสอบการกระจายของข้อมูลทั้ง 2 กลุ่ม เท่ากันหรือไม่

พิจารณาจากตาราง Levene's Test for Equality of Variances โดยสมมติฐานของการทดสอบการกระจายของข้อมูลทั้ง 2 กลุ่มเป็นแบบสองหาง (Two-tailed test) เขียนได้ ดังนี้

H_0 : การกระจายตัวของข้อมูลตัวอย่างกลุ่ม 1 = การกระจายตัวของข้อมูลตัวอย่างกลุ่ม 2

H_1 : การกระจายตัวของข้อมูลตัวอย่างกลุ่ม 1 $\neq$ การกระจายตัวของข้อมูลตัวอย่างกลุ่ม 2

เมื่อพิจารณาจากตารางช่อง Levene's Test for Equality of Variances พบว่าค่า Sig. เท่ากับ 0.498 ซึ่งมีค่ามากกว่า 0.05 สรุปว่า ยอมรับสมมติฐานหลัก (H_0) นั่นคือ การกระจายของข้อมูลตัวอย่างกลุ่ม 1 เท่ากับการกระจายของข้อมูลตัวอย่างกลุ่ม 2 ที่ระดับนัยสำคัญ 0.05 หรือ "Equal variances assumed"

(2) พิจารณาความแตกต่างค่าเฉลี่ยของทั้งสองกลุ่ม

จากขั้นตอนที่ 1 พบว่า การกระจายของข้อมูลทั้งสองกลุ่มไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05 หรือ "Equal variances assumed" ดังนั้นในขั้นตอนที่ 2 จะพิจารณาข้อมูลสถิติในส่วนของตาราง t-test for Equality of Means แลวบน นั่นคือ ค่า Sig. (2-tailed) เท่ากับ 0.002

จากผลลัพธ์ในภาพที่ 6.28 สามารถสรุปผลการวิเคราะห์ได้ 3 กรณีตามสมมติฐานในการทดสอบ ดังนี้

กรณีที่ 1 : ทดสอบสมมติฐานแบบสองหาง

H_0 : ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A = ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

H_1 : ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\neq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

จากตาราง t-test for Equality of Means แลวบน แสดง ค่า Sig. (2-tailed) เท่ากับ 0.002 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ จึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\neq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B ที่ระดับนัยสำคัญ 0.05

กรณีที่ 2 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางมากกว่า

H_0 : ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\leq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

H_1 : ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $>$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

จากตาราง t-test for Equality of Means แถวบน แสดง ค่า Sig. (2-tailed) เท่ากับ 0.002 เมื่อนำมาหารสองได้ค่าเท่ากับ 0.001 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ และค่า t มีค่าเป็นบวก ($t = 4.462$) ดังนั้นจึงสรุปได้ว่า ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $>$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B ที่ระดับนัยสำคัญ 0.05

กรณีที่ 3 : ทดสอบสมมติฐานแบบหางเดียว – ทิศทางน้อยกว่า

H_0 : ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\geq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

H_1 : ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $<$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B

จากตาราง t-test for Equality of Means แถวบน แสดง ค่า Sig. (2-tailed) เท่ากับ 0.002 เมื่อนำมาหารสองได้ค่าเท่ากับ 0.001 ซึ่งน้อยกว่าค่า $\alpha = 0.05$ และค่า t มีค่าเป็นบวก ($t = 4.462$) ดังนั้นจึงสรุปได้ว่า ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A $\geq$ ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B ที่ระดับนัยสำคัญ 0.05

(3) การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ค่าเฉลี่ยด้วยวิธี Independent-samples t-test

จากภาพที่ 6.28 สรุปได้ว่า ค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี A มีค่าไม่เท่ากับค่าเฉลี่ยปริมาณสารที่สกัดด้วยวิธี B ที่ระดับนัยสำคัญ 0.05 แสดงว่าผู้วิจัยไม่สามารถใช้วิธี B แทนที่วิธี A ในการสกัดสารจากตัวอย่างรำข้าวได้ เนื่องจากวิธี B ให้ค่าเฉลี่ยปริมาณสารสกัดน้อยกว่าวิธี A จากผลการวิเคราะห์ดังกล่าวสามารถเขียนผลการวิเคราะห์ในรูปแบบตารางได้ดังแสดงในตารางที่ 6.12

ตารางที่ 6.12 ปริมาณสารที่สกัดได้ (%Yield) ในตัวอย่างรำข้าวโดยใช้การสกัดด้วยวิธี A และวิธี B

วิธีการสกัดสาร	ปริมาณสารที่สกัดได้ (%)
วิธี A (n=4)	$26.48 \pm 0.13a$
วิธี B (n=7)	$26.01 \pm 0.18b$

^{a-b} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)



บทที่ 7

การวิเคราะห์ความแปรปรวน

วัตถุประสงค์ของผู้วิจัยส่วนใหญ่มักต้องการหาสาเหตุ หรือ หาเหตุและผลของสิ่งที่ศึกษา ในงานวิจัยด้านการพัฒนาผลิตภัณฑ์สามารถแบ่งการทดสอบค่าเฉลี่ยออกได้เป็น 3 กลุ่มใหญ่ตามจำนวนกลุ่มของตัวอย่างหรือประชากร ดังนี้

1) งานวิจัยที่เกี่ยวข้องกับการทดสอบค่าเฉลี่ยสำหรับ 1 กลุ่มตัวอย่าง เช่น พนักงานฝ่ายควบคุมคุณภาพต้องการทราบว่าวัตถุดิบที่เข้ามาส่งรอบนี้มีปริมาณโปรตีนเท่ากับ 25 เปอร์เซ็นต์ตามมาตรฐานที่กำหนดไว้หรือไม่ วิธีการทางสถิติที่นิยมใช้ทดสอบสมมติฐานงานวิจัย คือ One-sample t-test (รายละเอียดอ่านได้ในบทที่ 6 หัวข้อ 6.1)

2) งานวิจัยที่เกี่ยวข้องกับการทดสอบค่าเฉลี่ยสำหรับ 2 กลุ่มตัวอย่าง วิธีการทางสถิติที่นิยมใช้ทดสอบสมมติฐานงานวิจัย คือ Paired-samples t-test และ Independent-samples t-test (รายละเอียดอ่านได้ในบทที่ 6 หัวข้อ 6.2-6.4) ทั้งนี้การเลือกใช้ขึ้นอยู่กับความเป็นอิสระของข้อมูลในแต่ละกลุ่มตัวอย่าง เช่น

- ถ้านักวิจัยต้องการทราบว่าก่อนและหลังรับรู้ข้อมูลผลิตภัณฑ์ ผู้ทดสอบให้คะแนนความชอบ कै้กเสริมเส้นใยจากแป้งกล้วยดิบแตกต่างกันหรือไม่ กรณีนี้ข้อมูลในแต่ละกลุ่มตัวอย่างไม่เป็นอิสระต่อกัน เนื่องจากข้อมูลคะแนนความชอบก่อนและหลังทราบข้อมูลผลิตภัณฑ์มาจากผู้ทดสอบคนเดิม ดังนั้น จะต้องเลือกใช้วิธี Paired-samples t-test

- ถ้านักวิจัยต้องการทราบว่าผู้ทดสอบสองกลุ่ม (นิสิตปริญญาตรี และ นิสิตปริญญาโท) ให้คะแนนความชอบต่อผลิตภัณฑ์น้ำส้มสุตรลดน้ำตาลแตกต่างกันหรือไม่ กรณีนี้ข้อมูลในแต่ละกลุ่มตัวอย่างเป็นอิสระต่อกันหรือไม่มีความสัมพันธ์กัน ดังนั้นจะเลือกใช้วิธี Independent-samples t-test

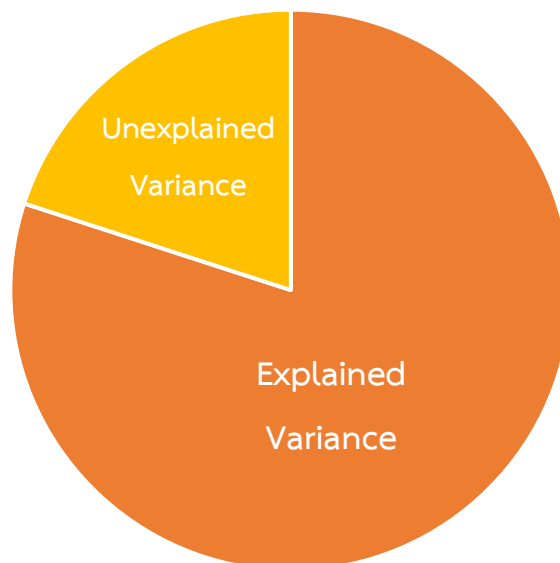
3) งานวิจัยที่เกี่ยวข้องกับการทดสอบค่าเฉลี่ยมากกว่า 2 กลุ่มตัวอย่างขึ้นไป เช่น นักวิจัยต้องการทราบว่าอุณหภูมิในการอบที่แตกต่างกัน 3 ระดับ (50, 60 และ 70 องศาเซลเซียส) มีผลทำให้ความชื้นของกล้วยตากแตกต่างกันหรือไม่ วิธีการทางสถิติที่นิยมใช้ทดสอบสมมติฐานงานวิจัย คือ การวิเคราะห์ความแปรปรวน (Analysis of Variance, ANOVA) และหากพบว่ามีความแตกต่างอย่างนัย 2 สิ่งทดลองแตกต่างกัน ขั้นตอนต่อไปที่ต้องวิเคราะห์คือการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple comparison test) ว่า สิ่งทดลองคู่ใดที่แตกต่างกันอย่างมีนัยสำคัญทางสถิติ

7.1 การวิเคราะห์ความแปรปรวน

การวิเคราะห์ความแปรปรวน (Analysis of Variance, ANOVA) เป็นวิธีการที่นิยมใช้ทดสอบค่าเฉลี่ยข้อมูลตั้งแต่ 3 กลุ่มขึ้นไป เป็นการทดสอบสมมติฐานเพื่อตรวจสอบว่าคุณลักษณะใดคุณลักษณะหนึ่งของข้อมูลตั้งแต่ 3 กลุ่มขึ้นไปนั้นมีความแตกต่างกันหรือไม่ และถ้าแตกต่างกันแล้วมีความแตกต่างกันอย่างไร โดยจะพิจารณาเปรียบเทียบจากค่าเฉลี่ยของคุณลักษณะนั้นๆ

การวิเคราะห์ความแปรปรวน นิยมเรียกว่า ANOVA เป็นการวิเคราะห์สถิติอย่างหนึ่งที่ใช้หาความสัมพันธ์ระหว่าง 2 ตัวแปร คือ (1) ตัวแปรต้น (Independent variable) บางครั้งเรียกว่า ปัจจัย (Factor) ซึ่งจะแสดงถึงการแบ่งกลุ่ม เป็นได้ทั้งข้อมูลเชิงปริมาณและเชิงคุณภาพ กับ (2) ตัวแปรตาม (Dependent variable) ซึ่งจะแสดงคุณลักษณะของกลุ่มและต้องเป็นข้อมูลเชิงปริมาณ

การวิเคราะห์ความแปรปรวนเป็นการศึกษาถึงสาเหตุที่ทำให้ค่าของตัวแปรตามในแต่ละกลุ่มแตกต่างกัน ซึ่งเรียกสาเหตุเหล่านี้ว่า แหล่งของความแปรปรวน (Source of Variation หรือ SOV) ในการวิเคราะห์ความแปรปรวนจะเป็นการเปรียบเทียบสัดส่วนของ ความแปรปรวนที่อธิบายได้ (Explained variance) และ ความแปรปรวนที่อธิบายไม่ได้ (Unexplained variance) ดังภาพที่ 7.1



ภาพที่ 7.1 การแบ่งสัดส่วนความแปรปรวน

7.2 ชนิดของความแปรปรวน

ชนิดของความแปรปรวนแบ่งออกได้เป็น 2 ชนิด คือ

7.2.1 ความแปรปรวนที่อธิบายได้ (Explained variance) เป็นความแปรปรวนที่สามารถบอกสาเหตุหรือแหล่งที่ก่อให้เกิดความแปรปรวนได้ เช่น ค่าความสว่าง L^* ของเยนสับปะรด 3 สิ่งทดลองแตกต่างกันเกิดจากปริมาณน้ำตาลที่ใส่เข้าไปในแต่ละสิ่งทดลองต่างกัน ซึ่งถือว่าเป็นความแปรปรวนที่สามารถอธิบายได้ว่าเกิดจากกลุ่มที่ต่างกัน หรือ ระดับของปัจจัยที่ต่างกัน ในกรณีนี้ปัจจัย คือ ปริมาณน้ำตาล

7.2.2 ความแปรปรวนที่อธิบายไม่ได้ (Unexplained variance หรือ Error variation) เป็นความแปรปรวนที่ไม่สามารถบอกสาเหตุหรือแหล่งที่ก่อให้เกิดความแปรปรวนได้ เช่น ตัวอย่างเผือกแผ่นทอดที่ทอดในสภาวะเดียวกันและใช้วัตถุดิบจากชุดเดียวกัน เมื่อสุ่มตัวอย่างมาวิเคราะห์ปริมาณไขมันพบว่าปริมาณไม่เท่ากันและมีความแตกต่างกันมาก ซึ่งถือว่าเป็นความแปรปรวนที่ไม่สามารถอธิบายได้ว่าเกิดจากสาเหตุใด

ในการวิเคราะห์ความแปรปรวนจะเป็นการเปรียบเทียบสัดส่วนของ ความแปรปรวนที่อธิบายได้ (Explained variance) ต่อ ความแปรปรวนที่อธิบายไม่ได้ (Unexplained variance) ดังนี้

$$\text{สัดส่วนความแปรปรวน} = \frac{\text{ความแปรปรวนที่อธิบายได้ (Explained variance)}}{\text{ความแปรปรวนที่อธิบายไม่ได้ (Unexplained variance)}}$$

การพิจารณาค่าสัดส่วนความแปรปรวน แบ่งเป็น 2 กรณี ดังนี้

กรณีที่ 1 ค่าสัดส่วนความแปรปรวนมีค่ามากกว่า 1 หมายถึง ความแปรปรวนที่เกิดขึ้นส่วนใหญ่เกิดจากกลุ่มที่ต่างกัน นั่นคือ ค่าเฉลี่ยของตัวแปรตามในแต่ละกลุ่มที่ต่างกันเกิดจากกลุ่มที่ต่างกัน หรือเกิดจากผลของตัวแปรต้นที่ระดับแตกต่างกัน เช่น ค่าการร่อนน้ำมันของโดนัทแป้งมันเทศ 4 สิ่งทดลองแตกต่างกันเกิดจากปริมาณแป้งมันเทศที่ใช้แทนที่แป้งสาลีในแต่ละสิ่งทดลองต่างกัน กรณีนี้ ตัวแปรตาม คือ ค่าการร่อนน้ำมัน (Oil uptake) และ ตัวแปรต้น คือ ปริมาณแป้งมันเทศที่ใช้แทนที่แป้งสาลี

กรณีที่ 2 ค่าสัดส่วนความแปรปรวนมีค่าน้อยกว่า 1 หมายถึง ความแปรปรวนที่เกิดขึ้นส่วนใหญ่เกิดจากความแตกต่างภายในกลุ่มนั้นๆ นั่นคือ ค่าเฉลี่ยของตัวแปรตามในแต่ละกลุ่มที่ต่างกันไม่ได้เกิดจากกลุ่มที่ต่างกัน หรือไม่ได้เกิดจากผลของตัวแปรต้นที่ระดับแตกต่างกัน

ในการตรวจสอบค่าสัดส่วนว่ามีค่ามากกว่าหรือน้อยกว่า 1 จะใช้วิธีทางสถิติที่เรียกว่า F-test

7.3 ประเภทของการวิเคราะห์ความแปรปรวน

การวิเคราะห์ความแปรปรวนสามารถแบ่งออกได้เป็น 3 ประเภทใหญ่ๆ ตามจำนวนตัวแปรต้นหรือปัจจัยที่ใช้ในการจำแนกกลุ่มข้อมูล ดังนี้

(1) การวิเคราะห์ความแปรปรวนแบบทางเดียว หรือ แบบมีปัจจัยเดียว (One-Way ANOVA)

เป็นการจำแนกข้อมูลด้วยตัวแปรต้นหรือปัจจัยเพียงปัจจัยเดียว หรือเป็นการวิเคราะห์ความแตกต่างกันของระดับต่างๆ ของปัจจัยที่สนใจนั่นเอง เช่น นักวิจัยคาดว่าปัจจัยที่ทำให้คะแนนความชอบน้ำผลไม้ต่างกันมีเพียงปัจจัยเดียว คือ อาชีพผู้ทดสอบ (นิสิต, อาจารย์ และ เจ้าหน้าที่) กรณีนี้จำแนกผู้ทดสอบโดยใช้อาชีพเป็นตัวจำแนก โดยแบ่งออกเป็น 3 กลุ่ม ได้แก่ นิสิต, อาจารย์ และ เจ้าหน้าที่ เรียกการวิเคราะห์ความแปรปรวนที่มีปัจจัยเดียวในการแบ่งกลุ่มนี้ว่า การวิเคราะห์ความแปรปรวนแบบทางเดียว หรือ One-Way ANOVA

วัตถุประสงค์ของการวิเคราะห์ความแปรปรวนแบบทางเดียว คือ การทดสอบความแตกต่างระหว่างค่าเฉลี่ยของประชากรหรือตัวอย่างที่รับปัจจัยที่ต่างระดับกันตั้งแต่ 3 ระดับขึ้นไป เช่นกรณีดังกล่าวข้างต้นคือการเปรียบเทียบค่าเฉลี่ยคะแนนความชอบน้ำผลไม้ของผู้ทดสอบ 3 กลุ่ม (มี 3 อาชีพ) เงื่อนไขของการทดสอบความแปรปรวนแบบทางเดียว คือ การสุ่มตัวอย่างแต่ละชุดต้องเป็นอิสระกัน, สุ่มตัวอย่างจากประชากรที่มีการแจกแจงแบบปกติ และ ความแปรปรวนของแต่ละประชากรต้องเท่ากัน

(2) การวิเคราะห์ความแปรปรวนแบบสองทาง หรือ แบบมีสองปัจจัย (Two-Way ANOVA) เป็น

การจำแนกข้อมูลหรือแบ่งข้อมูลด้วยตัวแปรต้นหรือปัจจัย 2 ปัจจัย และแต่ละปัจจัยมีหลายระดับ เช่น นักวิจัยคาดว่าปัจจัยที่ทำให้คะแนนความชอบน้ำผลไม้ของผู้ทดสอบต่างกันมี 2 ปัจจัย คือ เพศ (ชาย และ หญิง) และ อาชีพ (นิสิต, อาจารย์ และ เจ้าหน้าที่) กรณีนี้จำแนกผู้ทดสอบโดยใช้เพศและอาชีพเป็นตัวจำแนก จะเรียกการวิเคราะห์ความแปรปรวนที่มีสองปัจจัยนี้ว่า การวิเคราะห์ความแปรปรวนแบบสองทาง หรือ Two-Way ANOVA โดยปกติเมื่อมี 2 ปัจจัยที่มีผลต่อสิ่งที่ต้องการเปรียบเทียบ จะเรียกปัจจัยแรกว่า “ปัจจัย A” และเรียกปัจจัยที่สองว่า “ปัจจัย B” ซึ่งถือว่าทั้งสองปัจจัยเป็นปัจจัยหลัก (Main factor)

(3) การวิเคราะห์ความแปรปรวนแบบหลายทาง หรือ แบบมีหลายปัจจัย (Multi-Way ANOVA)

เป็นการจำแนกข้อมูลหรือแบ่งข้อมูลด้วยปัจจัยมากกว่า 2 ปัจจัยขึ้นไป และแต่ละปัจจัยมีหลายระดับ เช่น นักวิจัยคาดว่าปัจจัยที่ทำให้คะแนนความชอบน้ำผลไม้ของผู้ทดสอบต่างกันมี 3 ปัจจัย คือ เพศ (ชาย และ หญิง) อาชีพ (นิสิต, อาจารย์ และ เจ้าหน้าที่) และ ระดับการศึกษา (ต่ำกว่าปริญญาตรี และ สูงกว่าปริญญาตรี) กรณีนี้จำแนกผู้ทดสอบโดยใช้ 3 ปัจจัย ได้แก่ เพศ, อาชีพ และ ระดับการศึกษา จะเรียกการวิเคราะห์ความแปรปรวนที่มีมากกว่าสองปัจจัยนี้ว่า การวิเคราะห์ความแปรปรวนแบบหลายทาง หรือ Multi-Way ANOVA

ในหนังสือเล่มนี้จะกล่าวถึงเฉพาะการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) และ การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) เท่านั้น

7.4 ขั้นตอนการวิเคราะห์ความแปรปรวน

7.4.1 ข้อกำหนดสำหรับการวิเคราะห์ความแปรปรวน

(1) ข้อมูลที่นำมาวิเคราะห์ความแปรปรวน ได้แก่ ข้อมูลตัวแปรต้น (Independent variable) ซึ่งจะต้องแสดงถึงการแบ่งกลุ่ม (เป็นได้ทั้งข้อมูลเชิงปริมาณและเชิงคุณภาพ) กับ ข้อมูลตัวแปรตาม (Dependent variable) ซึ่งจะแสดงคุณลักษณะของกลุ่มและเป็นข้อมูลเชิงปริมาณที่สามารถนำมาคำนวณค่าเฉลี่ยได้ ได้แก่ ข้อมูลเชิงปริมาณที่ใช้สเกลอัตราส่วน หรือสเกลอันดับ

(2) ข้อมูลแต่ละกลุ่มย่อยต้องมาจากประชากรที่มีการกระจายหรือความแปรปรวนไม่แตกต่างกันทางสถิติ ($\sigma_1^2 = \sigma_2^2 = \sigma_3^2 \dots = \sigma_k^2$) นักวิจัยควรทำการตรวจสอบคุณสมบัติของข้อมูลตามข้อกำหนดนี้เนื่องจากการกระจายของประชากรแต่ละกลุ่มมีผลต่อความถูกต้องในการสรุปผลมากที่สุด วิธีการทางสถิติที่นิยมใช้ทดสอบ เช่น วิธี Bartlett-box F, วิธี Cochran และ วิธี Hartley

(3) ข้อมูลในแต่ละกลุ่มย่อยที่เลือกมาทดสอบควรมาจากประชากรที่มีการแจกแจงแบบปกติ ทั้งนี้ข้อกำหนดข้อนี้มีผลต่อการวิเคราะห์ความแปรปรวนน้อยมากถ้าข้อมูลไม่มีลักษณะเบ้ (Skewness) หรือมีความโด่ง (Kurtosis) อย่างเห็นได้ชัด อย่างไรก็ตามถ้าพบว่าข้อมูลที่ต้องการทดสอบมีลักษณะเบ้ นักวิจัยอาจแก้ไขโดยการเพิ่มขนาดหรือจำนวนตัวอย่างให้มากขึ้น

(4) ตัวอย่างที่เลือกมาแต่ละกลุ่มควรเป็นอิสระต่อกัน

7.4.2 ขั้นตอนการวิเคราะห์ความแปรปรวน

ในการวิเคราะห์ความแปรปรวนเพื่อทดสอบความแตกต่างของข้อมูลที่ได้รับปัจจัยที่ต่างระดับกัน จะทำการสร้างตารางการวิเคราะห์ความแปรปรวน (Analysis of Variance Table) หรือ เรียกว่า ANOVA Table โดยตัวแปรตามที่ใช้ในการวิเคราะห์ความแปรปรวนจะต้องเป็นตัวแปรเชิงปริมาณ และปัจจัยจะต้องเป็นตัวแปรแสดงกลุ่ม เช่น นักวิจัยต้องการทราบว่าคะแนนความชอบไอศกรีมขึ้นอยู่กั้อาชีพ (นิสิต, อาจารย์ และ เจ้าหน้าที่) หรือไม่ ในกรณีนี้ ปัจจัยหรือตัวแปรต้น คือ อาชีพ และ ตัวแปรตามคือ คะแนนความชอบ ซึ่งเขียนสมมติฐานในการทดสอบทางสถิติได้ ดังนี้

H_0 : ค่าเฉลี่ยคะแนนความชอบไอศกรีมไม่ขึ้นอยู่กั้อาชีพ

H_1 : ค่าเฉลี่ยคะแนนความชอบไอศกรีมขึ้นอยู่กั้อาชีพ

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

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

H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่ โดยที่ μ_i = ค่าเฉลี่ยคะแนนความชอบไอศกรีมของกลุ่มอาชีพที่ i ;
 $i = 1, 2$ และ 3 เมื่อสนใจเปรียบเทียบ 3 อาชีพ

เมื่อทำการตั้งสมมติฐาน วางแผนการทดลอง และเก็บรวบรวมข้อมูลแล้ว ก่อนจะทำการวิเคราะห์ความแปรปรวนจะต้องตรวจสอบเงื่อนไขของการทดสอบความแปรปรวนก่อน (ดังอธิบายในข้อ 7.4.1) ถ้าข้อมูลแต่ละกลุ่มย่อยมีการกระจายหรือความแปรปรวนไม่แตกต่างกันทางสถิติ และข้อมูลในแต่ละกลุ่มย่อยที่เลือกมาทดสอบมีการแจกแจงแบบปกติ ขั้นตอนต่อไปคือการทดสอบสมมติฐานงานวิจัยโดยพิจารณาค่าสถิติ F-test ที่คำนวณจากข้อมูลและแสดงในตารางการวิเคราะห์ความแปรปรวน (ANOVA table) โดยขั้นตอนการวิเคราะห์ความแปรปรวนสามารถแบ่งออกได้เป็น 3 ขั้นตอน ดังนี้

ขั้นที่ 1 การทดสอบสมมติฐานงานวิจัย

1.1 สมมติฐานทางสถิติที่ใช้ในการทดสอบ เขียนได้ ดังนี้

H_0 : ค่าเฉลี่ยของตัวอย่างทุกกลุ่มไม่แตกต่างกัน

H_1 : มีอย่างน้อย 2 กลุ่มตัวอย่างที่มีค่าเฉลี่ยแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\mu_1 = \mu_2 = \mu_3$ (กรณีมี 3 สิ่งทดลอง)

H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่; $i \neq j$

1.2 การตัดสินใจยอมรับหรือปฏิเสธสมมติฐานหลัก (H_0) พิจารณาจากค่าความน่าจะเป็นในการยอมรับสมมติฐานของ F หรือค่า Sig. ที่คำนวณได้จากตาราง ANOVA ดังนี้

- ถ้าค่าความน่าจะเป็นในการยอมรับสมมติฐานของ F หรือค่า Sig. ที่คำนวณได้จากตาราง ANOVA มีค่า **มากกว่าค่า α** ที่นักวิจัยกำหนดไว้ (เช่น ค่า $\alpha = 0.05$) จะตัดสินใจ **ยอมรับสมมติฐานหลัก (Accept H_0)** กรณีนี้ให้นักวิจัยหยุดการวิเคราะห์ ไม่ต้องทำต่อในขั้นตอนที่ 2 ให้สรุปผลการวิเคราะห์ในขั้นตอนที่ 3 ได้เลย

- ถ้าค่าความน่าจะเป็นในการยอมรับสมมติฐานของ F หรือค่า Sig. ที่คำนวณได้จากตาราง ANOVA มีค่า **น้อยกว่าหรือเท่ากับค่า α** ที่นักวิจัยกำหนดไว้ (เช่น ค่า $\alpha = 0.05$) จะตัดสินใจ **ปฏิเสธสมมติฐานหลัก (Reject H_0)** กรณีนี้ให้นักวิจัยทำการวิเคราะห์ต่อในขั้นตอนที่ 2

ขั้นที่ 2 การทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

จากขั้นตอนที่ 1 เมื่อนักวิจัยตัดสินใจปฏิเสธสมมติฐานหลัก (Reject H_0) นักวิจัยจะสามารถสรุปได้เพียงว่า มีตัวอย่างอย่างน้อย 2 กลุ่มที่มีค่าเฉลี่ยแตกต่างกันที่ระดับนัยสำคัญ 0.05 แต่ไม่ทราบว่าคู่ไหนที่แตกต่างกัน จึงต้องทำการเปรียบเทียบความแตกต่างของค่าเฉลี่ยโดยใช้วิธีการทดสอบพหุคูณ (Multiple comparison test) หรือบางครั้งเรียกว่า Post hoc tests โดยจะมีการคำนวณหาค่าผลต่างสัมบูรณ์ของค่าเฉลี่ยจากตัวอย่าง

ทุกคู่ที่เป็นไปได้ หรือ $|\bar{x}_i - \bar{x}_j|$ และนำค่าที่ได้ไปเปรียบเทียบกับค่าที่คำนวณได้จากวิธีการทางสถิติต่าง ๆ ภายใต้สมมติฐาน ดังนี้

H_0 : ค่าเฉลี่ยของตัวอย่าง 2 กลุ่มไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยของตัวอย่าง 2 กลุ่มแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\mu_i = \mu_j$

H_1 : $\mu_i \neq \mu_j ; i \neq j$

วิธีการทดสอบพหุคูณมีหลายวิธี แบ่งออกได้เป็น 2 กลุ่มตามความแปรปรวนของข้อมูลในแต่ละกลุ่ม (แตกต่าง/ไม่แตกต่าง) แต่ละวิธีมีข้อจำกัดแตกต่างกัน ไม่มีวิธีการใดที่เหมาะสมสำหรับทุกกรณี ในหนังสือเล่มนี้จะยกตัวอย่าง 3 วิธีที่นิยมใช้เมื่อความแปรปรวนของข้อมูลในแต่ละกลุ่มไม่แตกต่างกัน (Equal variance assumed) ได้แก่

(1) **วิธี LSD (Least Significant Difference)** เป็นวิธีการทดสอบความแตกต่างของค่าเฉลี่ย 2 กลุ่มโดยการเปรียบเทียบค่าของผลต่างเฉลี่ยกับค่าสถิติ LSD ที่เรียกว่า ผลต่างนัยสำคัญน้อยที่สุด (Least Significant Difference) มักใช้เปรียบเทียบค่าเฉลี่ยสำหรับสิ่งทดลองไม่เกิน 5 สิ่งทดลอง หากใช้กับสิ่งทดลองที่มากกว่านั้นอาจเกิดความผิดพลาดประเภทที่ 1 หรือ Type I error

(2) **วิธี Duncan** เป็นวิธีการทดสอบความแตกต่างของค่าเฉลี่ย 2 กลุ่มโดยการเปรียบเทียบค่าของผลต่างเฉลี่ยที่เรียงลำดับแล้วกับค่าสถิติ LSR ของ Duncan โดยอาจเรียกวิธีนี้ว่า Duncan's new multiple range test วิธีนี้มีโอกาสเกิดความผิดพลาดประเภทที่ 1 หรือ Type I error น้อยกว่าวิธี LSD

(3) **วิธี Tukey** เป็นวิธีการทดสอบความแตกต่างของค่าเฉลี่ย 2 กลุ่มโดยการเปรียบเทียบค่าของผลต่างเฉลี่ยกับค่าสถิติ HSD ของ Tukey โดยมีข้อกำหนดว่า จำนวนหรือขนาดตัวอย่างที่สุ่มมาแต่ละกลุ่มตัวอย่างต้องมีขนาดเท่ากัน วิธีการทดสอบ Tukey อาจเรียกอีกชื่อว่า Honestly Significant Difference หรือ HSD เมื่อเปรียบเทียบกับวิธี LSD และวิธี Duncan พบว่าวิธี Tukey มีโอกาสเกิดความผิดพลาดประเภทที่ 1 หรือ Type I error น้อยมาก แต่มีโอกาสเกิดความผิดพลาดประเภทที่ 2 หรือ Type II error มากกว่า

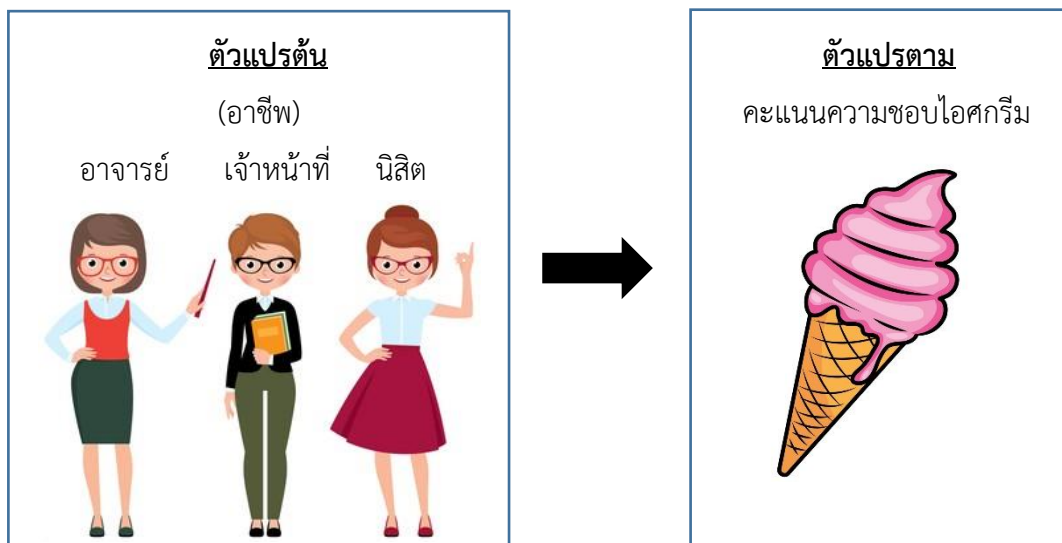
อาจกล่าวได้ว่า วิธีเปรียบเทียบค่าเฉลี่ยที่มีความไวหรืออำนาจการจำแนกกลุ่มสูงกว่าจะมีโอกาสเกิด Type I error มากกว่า ในทางตรงข้ามวิธีเปรียบเทียบค่าเฉลี่ยที่มีความไวต่ำกว่าจะเกิด Type II error มากกว่า

ขั้นที่ 3 สรุปผลและนำเสนอในรูปแบบตารางหรือกราฟแสดงให้เห็นว่าไม่แตกต่างหรือแตกต่างกันอย่างไร

รายละเอียดการนำเสนอผลการวิเคราะห์ในรูปแบบตารางแบ่งออกเป็น 3 ส่วนคือ ชื่อตาราง ข้อมูลในตาราง และรายละเอียดใต้ตาราง ซึ่งจะอธิบายรายละเอียดและแสดงตัวอย่างในหัวข้อถัดไป

7.5 การวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA)

การวิเคราะห์ความแปรปรวนแบบทางเดียว หรือ แบบมีปัจจัยเดียว (One-Way ANOVA) เป็นการจำแนกข้อมูลด้วยตัวแปรหรือปัจจัยเพียงปัจจัยเดียว เป็นการวิเคราะห์ความแตกต่างกันของระดับต่างๆ ของปัจจัยที่สนใจ เช่น นักวิจัยคาดว่าปัจจัยที่ทำให้คะแนนความชอบไอศกรีมต่างกันมีเพียงปัจจัยเดียว คือ อาชีพผู้ทดสอบ (นิสิต, อาจารย์ และ เจ้าหน้าที่) กรณีนี้ตัวแปรต้นที่ใช้จำแนกกลุ่มผู้ทดสอบ คือ อาชีพ แบ่งออกเป็น 3 กลุ่ม ได้แก่ นิสิต, อาจารย์ และ เจ้าหน้าที่ และตัวแปรตาม คือ คะแนนความชอบ



วัตถุประสงค์ของการวิเคราะห์ความแปรปรวนแบบทางเดียว คือ การทดสอบความแตกต่างระหว่างค่าเฉลี่ยของประชากรหรือตัวอย่างที่รับปัจจัยต่างระดับกันตั้งแต่ 3 ระดับขึ้นไป เช่นกรณีดังกล่าวข้างต้นคือการเปรียบเทียบค่าเฉลี่ยคะแนนความชอบไอศกรีมของผู้ทดสอบ 3 กลุ่ม (มี 3 อาชีพ) เงื่อนไขของการทดสอบความแปรปรวนแบบทางเดียว คือ การสุ่มตัวอย่างแต่ละชุดต้องเป็นอิสระกัน, สุ่มตัวอย่างจากประชากรที่มีการแจกแจงแบบปกติ และ ค่าความแปรปรวนของแต่ละประชากรต้องเท่ากัน

ในการวิเคราะห์ความแปรปรวนแบบทางเดียวจะมีการกำหนดตัวแบบ (Model) โดยกำหนดเป็นสมการทางคณิตศาสตร์ในรูปแบบเชิงเส้นตรง ซึ่งแบ่งออกได้เป็น 3 ประเภท คือ

(1) ตัวแบบผลกระทบชนิดคงที่ (Fixed Effect Model หรือ Model I)

ตัวแบบผลกระทบชนิดคงที่ เป็นตัวแบบชนิดที่มีปัจจัยคงที่ ซึ่งเป็นการนำค่าที่เป็นไปได้ของตัวแปรต้น (หรือระดับของปัจจัย) มาวิเคราะห์ทุกค่าหรือทุกระดับของปัจจัย ซึ่งหมายถึง ค่าของตัวแปรต้นเป็นค่าของประชากรเพราะนำมาหมดทุกค่า เช่น การทดสอบค่าเฉลี่ยเงินเดือนของทุกอาชีพในประเทศไทย ในกรณีนี้ตัวแปรต้น คือ อาชีพ โดยนำทุกอาชีพมาใช้ในการจำแนกกลุ่มเพื่อวิเคราะห์รายได้เฉลี่ยของแต่ละกลุ่ม ทำให้ตัวแปรต้น หรือ อาชีพ ที่นำมาใช้ในการวิเคราะห์คงที่ กล่าวคือ เป็นอาชีพเหล่านี้ตลอดไม่ว่าจะทำการเก็บข้อมูลหรือสำรวจมาเมื่อใด

(2) ตัวแบบผลกระทบชนิดสุ่ม (Random Effect Model หรือ Model II)

ตัวแบบผลกระทบชนิดสุ่ม เป็นตัวแบบชนิดที่มีปัจจัยไม่คงที่ ซึ่งเป็นการนำค่าที่เป็นไปได้ของตัวแปรต้น (หรือระดับของปัจจัย) มาวิเคราะห์บางค่าโดยเลือกมาอย่างสุ่ม ซึ่งหมายถึง ค่าของตัวแปรต้นเป็นค่าของตัวอย่างเพราะไม่ได้นำมาหมดทุกค่า เช่น การทดสอบค่าเฉลี่ยเงินเดือนของอาชีพในประเทศไทย ในกรณีนี้ตัวแปรต้น คือ อาชีพ โดยเลือกมาบางอาชีพเพื่อใช้ในการจำแนกกลุ่มสำหรับวิเคราะห์รายได้เฉลี่ยของแต่ละกลุ่ม ทำให้ตัวแปรต้น หรือ อาชีพ ที่นำมาใช้ในการวิเคราะห์ไม่คงที่ กล่าวคือ อาชีพที่นำมาวิเคราะห์จะเปลี่ยนแปลงไปสำหรับการวิเคราะห์หรือเก็บข้อมูลแต่ละครั้งขึ้นอยู่กับว่าในครั้งนั้นๆ สุ่มเลือกอาชีพใด

(3) ตัวแบบผลกระทบชนิดผสม (Mixed Effect Model หรือ Model III)

ตัวแบบผลกระทบชนิดผสม เป็นตัวแบบชนิดที่มีปัจจัยผสม ซึ่งจะเกิดขึ้นในการวิเคราะห์ความแปรปรวนตั้งแต่ 2 ทางขึ้นไปที่มีการแบ่งกลุ่มข้อมูลโดยตัวแปรต้นหรือปัจจัยมากกว่า 1 ปัจจัยนั่นเอง ทั้งนี้ปัจจัยที่ใช้ในการจำแนกกลุ่มนั้น บางปัจจัยเป็นปัจจัยคงที่ และบางปัจจัยเป็นแบบสุ่ม เช่น การทดสอบค่าเฉลี่ยเงินเดือนของอาชีพในประเทศไทยโดยจำแนกตามประเภทที่อยู่อาศัย (บ้าน, คอนโดมิเนียม และ หอพัก) ในกรณีนี้ประเภทที่อยู่อาศัย เป็นปัจจัยคงที่ และอาชีพที่สุ่มมาบางอาชีพ เป็นปัจจัยสุ่มหรือปัจจัยไม่คงที่

7.6 คำสั่ง SPSS ในการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA)

คำสั่งที่ใช้ในการเปรียบเทียบค่าเฉลี่ยของหลายประชากรว่าแตกต่างกันเนื่องจากอิทธิพลของปัจจัยใดปัจจัยหนึ่ง ทำได้ด้วยคำสั่ง 2 แบบ

แบบที่ 1 Analyze ➤ Compare Means ➤ One-Way ANOVA...

แบบที่ 2 Analyze ➤ General Linear Model ➤ Univariate...

ในหนังสือเล่มนี้จะอธิบายเฉพาะแบบที่ 2 เนื่องจากสามารถใช้วิเคราะห์ ANOVA ได้ทุกประเภท

7.6.1 ตัวอย่างข้อมูล

นักวิจัยพัฒนาผลิตภัณฑ์ชาสมุนไพรสมุนไพรไทย โดยนักวิจัยคาดว่าคะแนนความชอบโดยรวมของผลิตภัณฑ์นี้ขึ้นอยู่กับอาชีพของผู้ทดสอบ จึงได้ทำการประเมินความชอบผลิตภัณฑ์ชาสมุนไพรกับกลุ่มผู้ทดสอบ 3 กลุ่มแบ่งตามอาชีพ ได้แก่ นิสิต, อาจารย์ และ เจ้าหน้าที่ กลุ่มละ 15 คน โดยใช้สเกล 9-point hedonic scale (1=ชอบน้อยที่สุด และ 9=ชอบมากที่สุด) ได้ผลคะแนนความชอบโดยรวมดังแสดงในตารางที่ 7.1

ตารางที่ 7.1 คะแนนความชอบโดยรวมของผลิตภัณฑ์ชาสมุนไพรสชาไทยประเมินโดยผู้ทดสอบต่างอาชีพกัน

ผู้ทดสอบคนที่	คะแนนความชอบโดยรวม		
	กลุ่มนิสิต	กลุ่มอาจารย์	กลุ่มเจ้าหน้าที่
1	5	7	6
2	6	7	6
3	4	7	6
4	5	7	6
5	4	6	8
6	4	6	8
7	4	6	8
8	5	6	8
9	5	6	7
10	5	7	7
11	6	7	8
12	5	7	7
13	5	7	8
14	4	7	7
15	6	8	7

หมายเหตุ ในการประเมินความชอบทางด้านประสาทสัมผัสควรใช้ผู้ทดสอบไม่น้อยกว่า 50 คน

7.6.2 การเขียนสมมติฐานทางสถิติสำหรับการวิเคราะห์ความแปรปรวนแบบทางเดียว

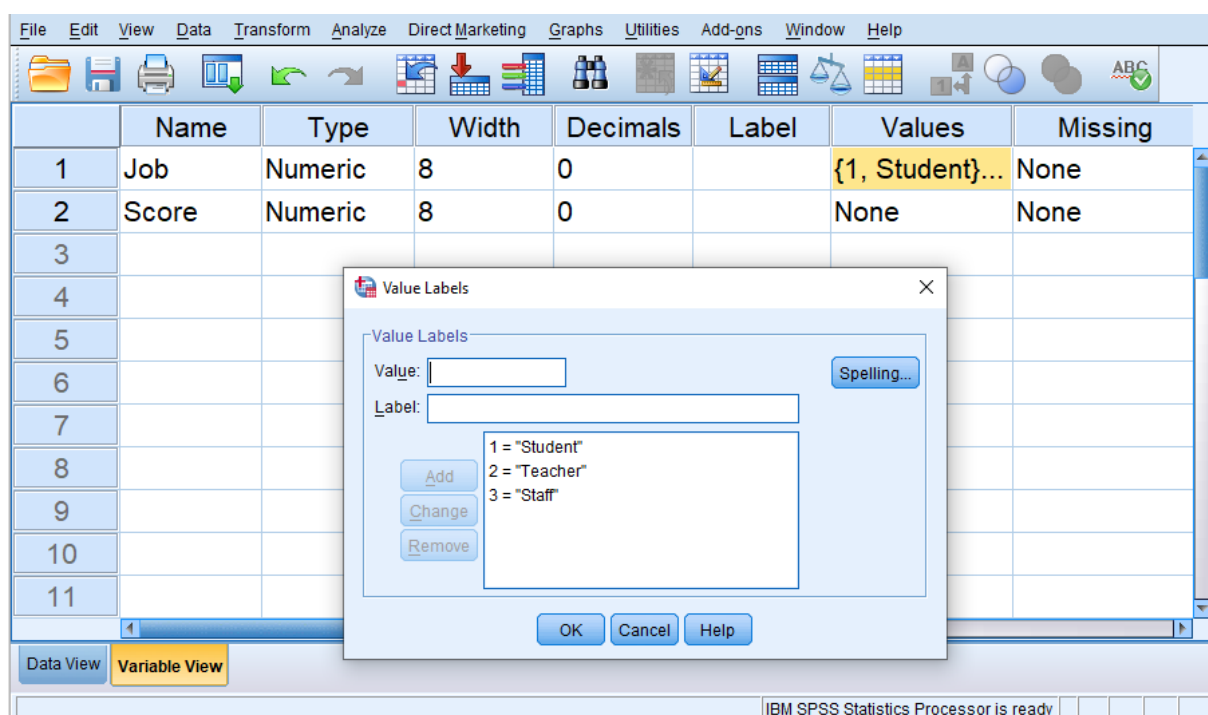
จากตัวอย่าง นักวิจัยสามารถเขียนสมมติฐานสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ ดังนี้

- H_0 : ค่าเฉลี่ยคะแนนความชอบชาสมุนไพรสชาไทยไม่ขึ้นอยู่กับอาชีพของผู้ทดสอบ
 H_1 : ค่าเฉลี่ยคะแนนความชอบชาสมุนไพรสชาไทยขึ้นอยู่กับอาชีพของผู้ทดสอบ
 หรือ H_0 : ค่าเฉลี่ยคะแนนความชอบชาสมุนไพรสชาไทยทั้ง 3 อาชีพไม่แตกต่างกัน
 H_1 : ค่าเฉลี่ยคะแนนความชอบชาสมุนไพรสชาไทยอย่างน้อย 2 อาชีพแตกต่างกัน
 หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้
 H_0 : $\mu_1 = \mu_2 = \mu_3$
 H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่; $i \neq j$ โดยที่ μ_i = คะแนนความชอบรวมของอาชีพที่ i

7.6.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ความแปรปรวนแบบทางเดียว

(1) เปิดโปรแกรม SPSS ในกรณีนี้การป้อนข้อมูลคะแนนความชอบของผู้ทดสอบแต่ละกลุ่ม จะต้องใช้ตัวแปร var จำนวน 2 ตัวแปรสำหรับป้อนข้อมูล 2 ตัว ได้แก่ กลุ่มผู้ทดสอบ และ คะแนนความชอบ ดังนั้นกำหนดชื่อตัวแปรเป็น **Job** และ **Score** ในหน้าต่าง Variable View ดังภาพที่ 7.2 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

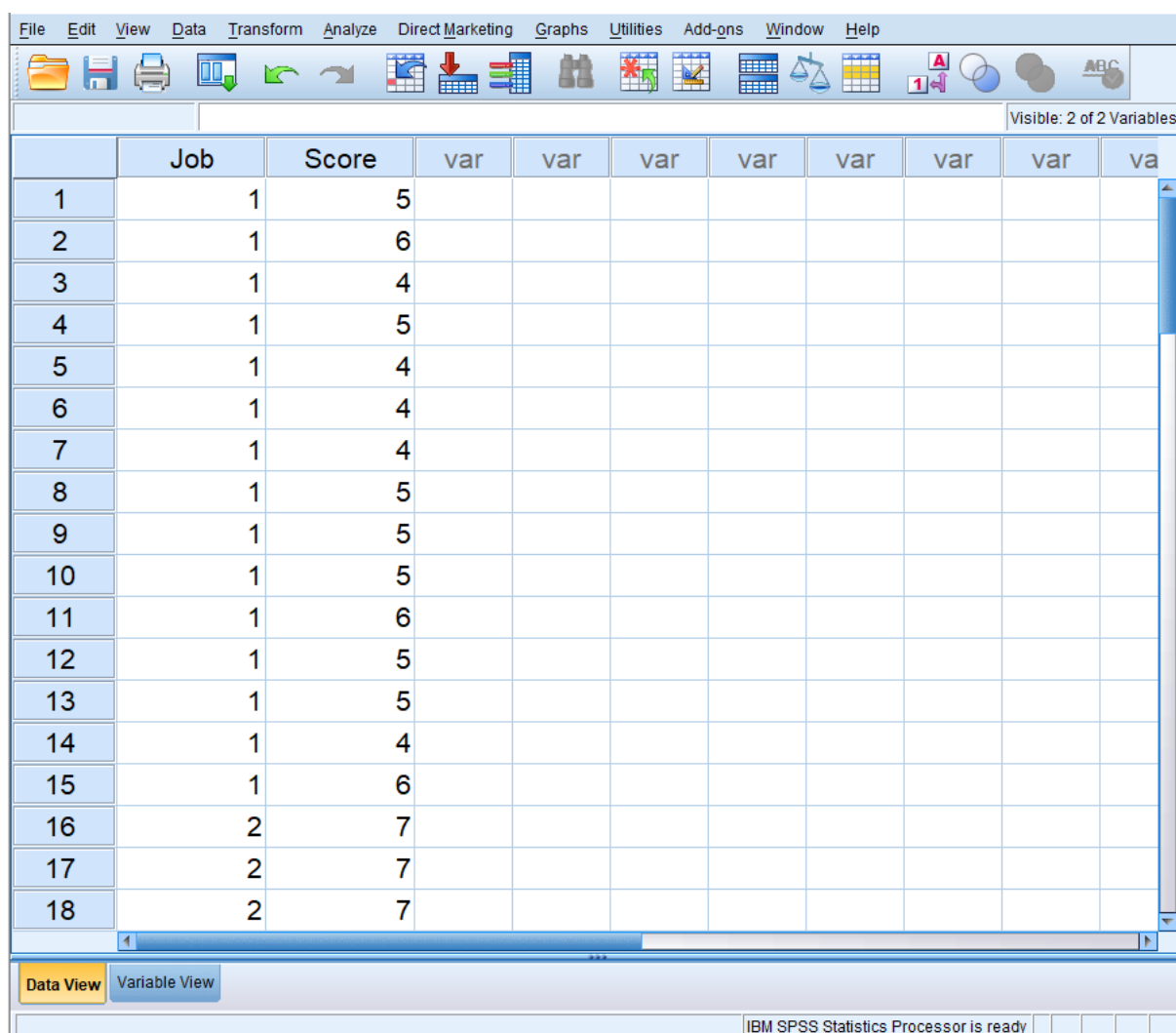
ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
กลุ่มผู้ทดสอบ	Job	ไม่ต้องกำหนด	0	1 = Student 2 = Teacher 3 = Staff
คะแนนความชอบ	Score	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด



ภาพที่ 7.2 กำหนดรายละเอียดตัวแปรเป็น Job และ Score ในหน้าต่าง Variable View

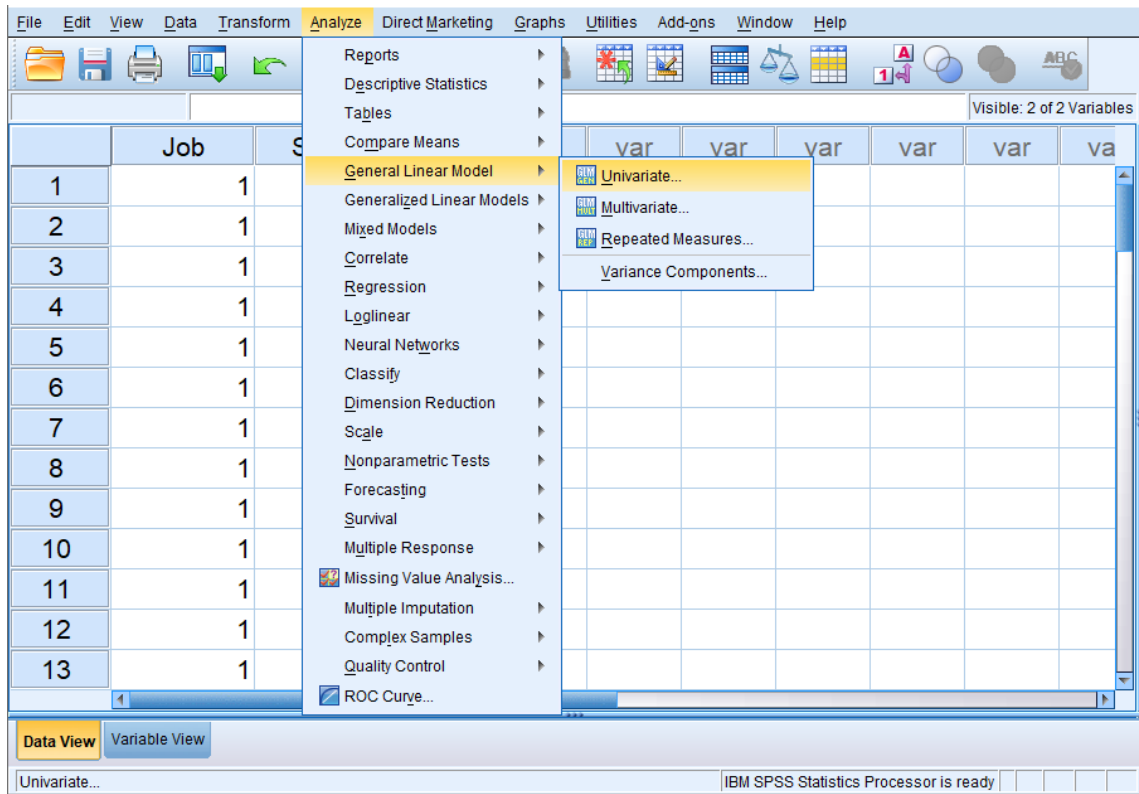
(วิธีการกำหนดค่า Values ให้ตัวแปร สามารถอ่านเพิ่มเติมได้ในบทที่ 3)

(2) ป้อนข้อมูลคะแนนความชอบเรียงตามกลุ่มอาชีพในหน้าต่าง Data View ดังภาพที่ 7.3

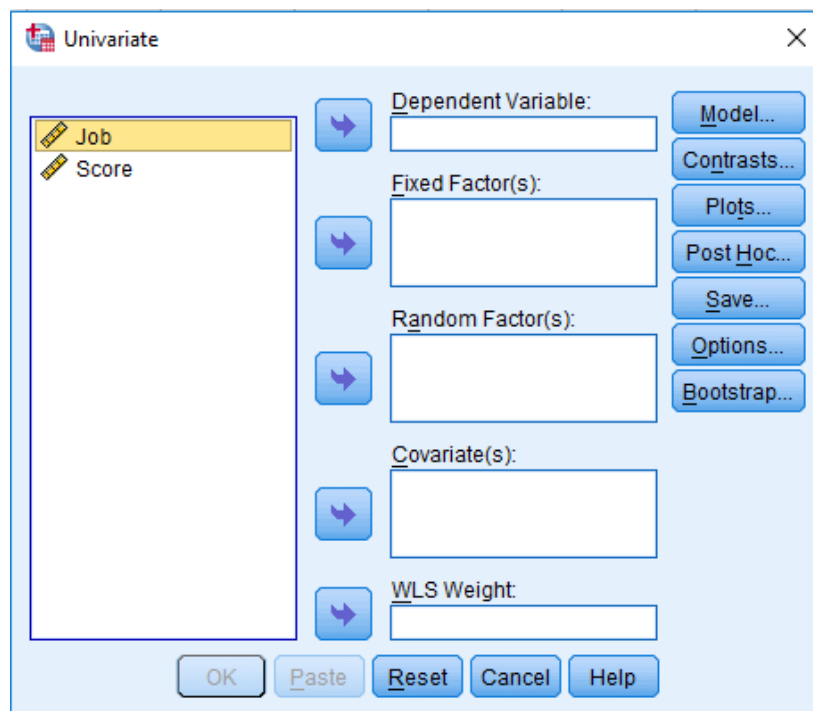


ภาพที่ 7.3 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze ► General Linear Model ► Univariate...** ดังภาพที่ 7.4 จะปรากฏหน้าต่างดังภาพที่ 7.5

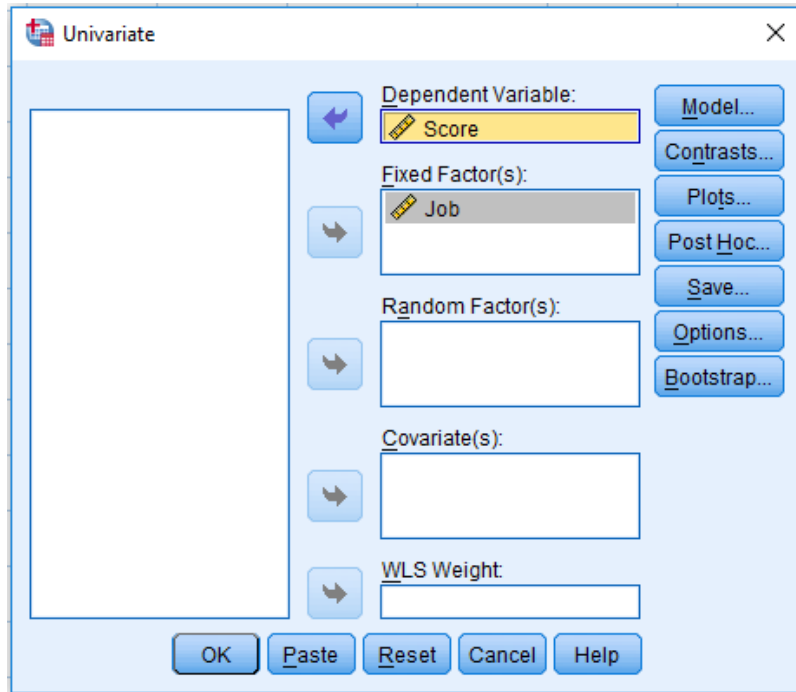


ภาพที่ 7.4 เมนูหน้าจอคำสั่ง Analyze ➤ General Linear Model ➤ Univariate...



ภาพที่ 7.5 หน้าจอคำสั่ง Univariate

(4) เลือกตัวแปร **Score** (ตัวแปรตาม) ใส่ในกล่อง Dependent Variable และตัวแปร **Job** (ตัวแปรต้น) ใส่ในกล่อง Fixed Factor(s) จะได้ดังภาพที่ 7.6

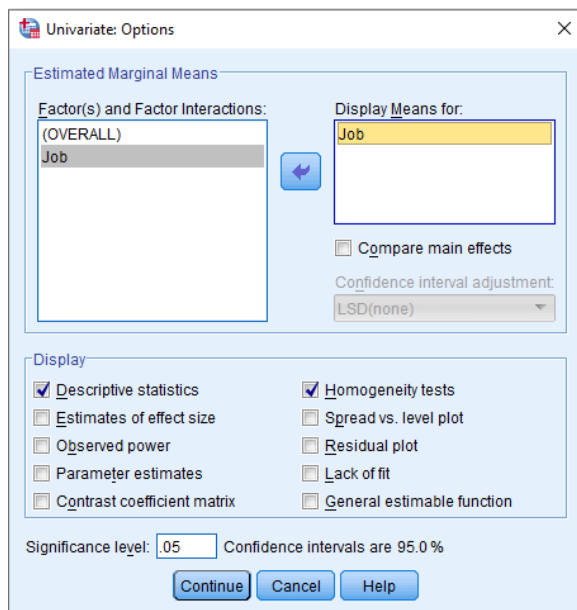


ภาพที่ 7.6 เลือกตัวแปร Score ใส่ในกล่อง Dependent Variable และตัวแปร Job ใส่ในกล่อง Fixed Factor(s)

ความหมายของ Fixed Factor และ Random Factor

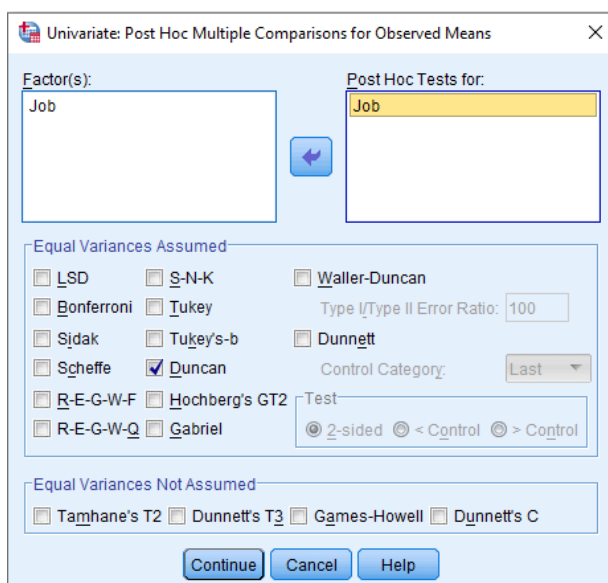
ในการวิเคราะห์ความแปรปรวนนั้น นักวิจัยต้องการศึกษาปัจจัยที่คาดว่าจะทำให้ตัวแปรตาม (ตัวแปรเชิงปริมาณ) มีค่าเฉลี่ยต่างกัน เช่น กรณีนี้นักวิจัยคาดว่าคะแนนความชอบผลิตภัณฑ์จะแตกต่างกันไปในแต่ละกลุ่มอาชีพ ในที่นี้ กลุ่มอาชีพ คือ ปัจจัยที่สนใจ สมมติอาชีพที่มีอยู่ในประเทศไทยมี 30 อาชีพ (เท่ากับตัวแปรต้นหรือปัจจัยที่ใช้ในการจำแนกกลุ่ม มี 30 กลุ่ม หรือ 30 ระดับ) แต่นักวิจัยสนใจศึกษาเพียงแค่ 3 อาชีพ ได้แก่ นิสิต, อาจารย์ และ เจ้าหน้าที่ (นักวิจัยระบุหรือเลือกเอง) กรณีนี้จะเรียกอาชีพที่ต้องการเปรียบเทียบเป็น ปัจจัยคงที่ (Fixed Factor) แต่ถ้านักวิจัยสุ่มมาเปรียบเทียบ 3 อาชีพ จาก 30 อาชีพ กรณีนี้จะถือว่าอาชีพเป็น ปัจจัยสุ่ม (Random Factor) นอกจากนี้ถ้าเปรียบเทียบทั้ง 30 อาชีพ ก็ถือว่าอาชีพเป็น ปัจจัยคงที่ (Fixed Factor) เช่นกัน

(5) เลือกคำสั่ง **Options...** จะได้หน้าจอตั้งภาพที่ 7.7 ให้เลือก **ตัวแปร Job** จากกล่องซ้ายมือไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลในแต่ละกลุ่มอาชีพ) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละกลุ่มอาชีพ) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 7.7 หน้าจอคำสั่ง Univariate : Options

(6) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 7.8 ให้กดเลือก **ตัวแปร Job** จากกล่องซ้ายมือไปใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธีอื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 7.9



ภาพที่ 7.8 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

Univariate Analysis of Variance

Between-Subjects Factors			
Job		Value Label	N
1		Student	15
2		Teacher	15
3		Staff	15

Descriptive Statistics			
Dependent Variable: Score			
Job	Mean	Std. Deviation	N
Student	4.87	.743	15
Teacher	6.73	.594	15
Staff	7.13	.834	15
Total	6.24	1.228	45

Levene's Test of Equality of Error Variances ^a			
Dependent Variable: Score			
F	df1	df2	Sig.
1.004	2	42	.375

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Job

Tests of Between-Subjects Effects					
Dependent Variable: Score					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	43.911 ^a	2	21.956	41.167	.000
Intercept	1754.689	1	1754.689	3290.042	.000
Job	43.911	2	21.956	41.167	.000
Error	22.400	42	.533		
Total	1821.000	45			
Corrected Total	66.311	44			

a. R Squared = .662 (Adjusted R Squared = .646)

Estimated Marginal Means

Job				
Dependent Variable: Score				
Job	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Student	4.867	.189	4.486	5.247
Teacher	6.733	.189	6.353	7.114
Staff	7.133	.189	6.753	7.514

Post Hoc Tests

Job

Homogeneous Subsets

Score			
Duncan ^{a, b}			
Job	N	Subset	
		1	2
Student	15	4.87	
Teacher	15		6.73
Staff	15		7.13
Sig.		1.000	.141

Means for groups in homogeneous subsets are displayed. Based on observed means. The error term is Mean Square(Error) = .533.

a. Uses Harmonic Mean Sample Size = 15.000.

b. Alpha = .05.

ภาพที่ 7.9 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนแบบทางเดียวด้วยคำสั่ง Univariate

7.6.4 การอ่านผลลัพธ์จากการวิเคราะห์ความแปรปรวนแบบทางเดียว

จากผลการวิเคราะห์ในภาพที่ 7.9 แสดงผลลัพธ์ได้ 6 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Between-Subjects Factors เป็นส่วนที่แสดงรายละเอียดของข้อมูลที่ป้อน ดังนี้

Job	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 3 กลุ่ม ได้แก่ 1, 2 และ 3
Value Label	คือ	ค่าความหมายที่แท้จริงของแต่ละกลุ่ม
		ในที่นี้ 1 คือ Student, 2 คือ Teacher และ 3 คือ Staff
N	คือ	จำนวนข้อมูลในแต่ละกลุ่ม ในที่นี้แต่ละกลุ่มมีจำนวนข้อมูลเท่ากัน คือ 15

ส่วนที่ 2 ตาราง Descriptive Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Job	คือ ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 3 กลุ่ม ได้แก่ Student, Teacher, Staff
Mean	คือ ค่าเฉลี่ยคะแนนความชอบของแต่ละกลุ่มจำแนกโดยอาชีพ
Std. Deviation	คือ ค่าส่วนเบี่ยงเบนมาตรฐานของคะแนนความชอบ แสดงการกระจายของข้อมูล
N	คือ จำนวนข้อมูลในแต่ละกลุ่ม ในที่นี้แต่ละกลุ่มมีจำนวนข้อมูลเท่ากัน คือ 15

ส่วนที่ 3 ตาราง Levene's Test of Equality of Error Variances ใช้พิจารณาการกระจายของข้อมูลหรือความแปรปรวนของข้อมูลแต่ละกลุ่ม อธิบายได้ ดังนี้

F, df1, df2	คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในกรณีนี้การทดสอบสมมติฐานเป็นแบบสองหาง (Two-tailed test) เขียนได้ดังนี้ H_0 : ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวไม่แตกต่างกัน H_1 : ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวแตกต่างกัน ค่า Sig. มีค่าเท่ากับ 0.375 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05 (Equal variance assumed)

ส่วนที่ 4 ตาราง Tests of Between-Subjects Effects เป็นส่วนที่แสดงค่าสถิติต่างๆ ของตารางวิเคราะห์ความแปรปรวน เพื่อใช้ในการทดสอบสมมติฐานทางสถิติด้วยค่าสถิติ F หรือค่าความน่าจะเป็น Sig. เพื่อใช้ทดสอบสมมติฐาน อธิบายได้ ดังนี้

Type III Sum of Squares, df, Mean Square, F	คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในที่นี้จะพิจารณาเฉพาะค่า Sig. ของตัวแปร Job เท่านั้น ตามสมมติฐานที่เขียนไว้ ดังนี้ H_0 : ค่าเฉลี่ยคะแนนความชอบขานมไข่มุกรสชาไทยทั้ง 3 อาชีพไม่ต่างกัน H_1 : ค่าเฉลี่ยคะแนนความชอบขานมไข่มุกรสชาไทยอย่างน้อย 2 อาชีพต่างกัน ค่า Sig. มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบขานมไข่มุกรสชาไทยอย่างน้อย 2 อาชีพต่างกันที่ระดับนัยสำคัญ 0.05

ส่วนที่ 5 ตาราง Estimated Marginal Means เป็นส่วนที่แสดงค่าเฉลี่ยโดยประมาณของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Job	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 3 กลุ่ม ได้แก่ Student, Teacher, Staff
Mean	คือ	ค่าเฉลี่ยคะแนนความชอบของแต่ละกลุ่มโดยประมาณ
Std. Error	คือ	ค่าความคลาดเคลื่อนมาตรฐานในแต่ละกลุ่ม
95% Confidence Interval	คือ	ค่าที่แสดงขอบเขตบนล่างของค่าเฉลี่ยในแต่ละกลุ่มในช่วงความเชื่อมั่น 95%

ส่วนที่ 6 ตาราง Post Hoc Tests เป็นส่วนที่แสดงค่าสถิติสำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Test) ซึ่งส่วนนี้จะนำมาใช้พิจารณาต่อเมื่อในส่วนที่ 4 ตาราง Tests of Between-Subjects Effects มีการปฏิเสธสมมติฐานหลัก (H_0) อธิบายได้ ดังนี้

Job	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 3 กลุ่ม ได้แก่ Student, Teacher, Staff
N	คือ	จำนวนข้อมูลในแต่ละกลุ่ม ในที่นี้แต่ละกลุ่มมีจำนวนข้อมูลเท่ากัน คือ 15
Subset	คือ	แสดงการแบ่งกลุ่มให้พิจารณารายละเอียดของ Subset ดังนี้ (1) มี 2 Subset แสดงว่ามี 2 กลุ่มที่แตกต่างกัน โดย Subset ที่ 1 มีสมาชิก 1 กลุ่มตัวอย่าง คือ Student และ Subset ที่ 2 มีสมาชิก 2 กลุ่มตัวอย่าง คือ Teacher และ Staff แสดงให้เห็นว่าค่าคะแนนความชอบเฉลี่ยของอาจารย์และเจ้าหน้าที่ไม่แตกต่างกันเนื่องจากอยู่ใน Subset เดียวกัน (2) ตัวเลขที่แสดงในแต่ละ Subset คือ ค่าเฉลี่ยคะแนนความชอบของอาชีพนั้นๆ เช่น Subset ที่ 1 แสดงค่าเฉลี่ยคะแนนความชอบของกลุ่ม Student มีค่าเท่ากับ 4.87 (3) แต่ละ Subset จะใช้ตัวอักษรภาษาอังกฤษพิมพ์เล็กในการบ่งชี้กลุ่ม โดยจะนำตัวอักษรนี้ไปแสดงหลัง ค่าเฉลี่ย $\pm$ ค่าส่วนเบี่ยงเบนมาตรฐาน เช่น $4.87 \pm 0.74b$ กรณีนี้มี 2 Subset จึงใช้ตัวอักษร a และ b บ่งชี้สมาชิกที่อยู่ใน Subset ที่ 2 และ 1 ตามลำดับ ซึ่งเป็นการเรียงลำดับอักษรจากค่าเฉลี่ยมากไปน้อย (นักวิจัยสามารถเลือกใช้อักษรเรียงลำดับ Subset ที่ 1 และ 2 ได้ ซึ่งจะเป็นการเรียงลำดับอักษรจากค่าเฉลี่ยน้อยไปมาก) (4) สมาชิกที่อยู่ใน Subset เดียวกันจะใช้อักษรตัวเดียวกัน

7.6.5 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ความแปรปรวนแบบทางเดียว

จากภาพที่ 7.9 ในตาราง Tests of Between-Subjects Effects มีการปฏิเสธสมมติฐานหลัก (H_0) จึงสรุปได้ว่า ค่าเฉลี่ยคะแนนความชอบชาไข่มุกรสชาติไทยอย่างน้อย 2 อาชีพแตกต่างกันที่ระดับนัยสำคัญ 0.05 จึงทำการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Test) ด้วยวิธี Duncan แล้วพบว่าค่าเฉลี่ยคะแนนความชอบของอาจารย์และเจ้าหน้าที่ไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05 โดยในการเขียนผลการวิเคราะห์ในรูปแบบตารางจะนำข้อมูลจากส่วนที่ 2 ตาราง Descriptive Statistics และส่วนที่ 6 ตาราง Post Hoc Tests มาใช้ในการเขียนผล ดังแสดงในตารางที่ 7.2

ตารางที่ 7.2 ค่าเฉลี่ยคะแนนความชอบชาไข่มุกรสชาติไทยของผู้ทดสอบที่มีอาชีพต่างกัน ประเมินโดยใช้สเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด)

อาชีพผู้ทดสอบ	คะแนนความชอบ
นิสิต	4.9 ± 0.7b
อาจารย์	6.7 ± 0.6a
เจ้าหน้าที่	7.1 ± 0.8a

^{a-b} ตัวอักษรที่ต่างกันแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

กรณีสมมติ หากวิเคราะห์ความแปรปรวนแล้วข้อมูลในตาราง Tests of Between-Subjects Effects พบว่าค่า Sig. มากกว่า 0.05 กรณีนี้ต้อง ยอมรับสมมติฐานหลัก (H_0) จะสรุปได้ว่า ค่าเฉลี่ยคะแนนความชอบชาไข่มุกรสชาติไทยของ 3 อาชีพไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05 นักวิจัยไม่ต้องพิจารณาผลการวิเคราะห์ในส่วนที่ 6 ตาราง Post Hoc Tests และสามารถเขียนสรุปผลในรูปแบบตารางได้ดังแสดงในตารางที่ 7.3 โดยใช้ตัวอักษร ns แสดงความไม่แตกต่างทางสถิติหลังซื้อตัวแปรตามที่น่ามาวิเคราะห์ความแปรปรวน

ตารางที่ 7.3 ค่าเฉลี่ยคะแนนความชอบชาไข่มุกรสชาติไทยของผู้ทดสอบที่มีอาชีพต่างกัน ประเมินโดยใช้สเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด)

อาชีพผู้ทดสอบ	คะแนนความชอบ ^{ns}
นิสิต	4.9 ± 0.7
อาจารย์	6.7 ± 0.6
เจ้าหน้าที่	7.1 ± 0.8

^{ns} ค่าเฉลี่ยไม่มีความแตกต่างกันทางสถิติ ($p > 0.05$)

7.7 กรณีศึกษาการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA)

7.7.1 ตัวอย่างข้อมูล

นิสิตปริญญาโทคนหนึ่งต้องการศึกษาผลของการใช้แป้งเผือกแทนที่แป้งสาลีบางส่วนต่อคุณภาพด้านสีของโดนัทอบ จึงได้ทำการแทนที่แป้งสาลีด้วยแป้งเผือก 4 ระดับ คือ 0, 25, 50 และ 75% จากนั้นนำโดนัทที่ได้ไปวิเคราะห์ค่าคุณภาพด้านสีด้วยระบบ CIE $L^*a^*b^*$ ด้วยเครื่อง Chroma meter โดยค่า L^* แสดงถึงค่าความสว่าง, ค่า a^* แสดงถึงค่าความเป็นสีแดง (เครื่องหมาย +) และสีเขียว (เครื่องหมาย -) และ ค่า b^* แสดงถึงค่าความเป็นสีเหลือง (เครื่องหมาย +) และสีน้ำเงิน (เครื่องหมาย -) ได้ผลการวิเคราะห์แสดงดังตารางที่ 7.4

ตารางที่ 7.4 ค่าคุณภาพด้านสีของโดนัทอบเมื่อใช้แป้งเผือกแทนที่แป้งสาลีที่ระดับต่างกัน

ปริมาณแป้งเผือกที่ใช้แทนที่แป้งสาลีบางส่วน (%)	วัดค่าซ้ำที่	ค่าสี L^*	ค่าสี a^*	ค่าสี b^*
0	1	57.55	8.26	37.01
	2	57.89	8.36	37.23
	3	57.64	8.54	37.22
25	1	50.65	14.53	34.38
	2	50.55	14.55	34.10
	3	50.45	14.62	34.28
50	1	41.23	14.65	28.41
	2	41.65	14.77	28.33
	3	41.87	14.75	28.45
75	1	38.78	16.89	26.97
	2	38.84	16.55	26.87
	3	38.69	16.73	26.65

จากข้อมูลดังกล่าว พิจารณารายละเอียดของข้อมูลได้ ดังนี้

- (1) ตัวแปรต้นหรือปัจจัยที่นิสิตปริญญาโทสนใจ มี 1 ปัจจัย คือ ระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี ดังนั้น จึงเป็นการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA)
- (2) ระดับของตัวแปรต้น เท่ากับ 4 ระดับ (0, 25, 50 และ 75%) ดังนั้น จึงมี 4 สิ่งทดลอง
- (3) ตัวแปรตาม มี 3 ตัว คือ ค่าสี L^* , a^* และ b^*
- (4) ต้องวิเคราะห์ One-Way ANOVA จำนวน 3 ครั้ง โดยวิเคราะห์ตัวแปรตามทีละตัว

7.7.2 การเขียนสมมติฐานทางสถิติสำหรับวิเคราะห์ความแปรปรวนแบบทางเดียว

นิสิตปริญญาโทสามารถเขียนสมมติฐานสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ 3 ข้อตามจำนวนตัวแปรตามที่ต้องการวิเคราะห์ ดังนี้

สมมติฐานข้อที่ 1 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย L^*

- H_0 : ค่าเฉลี่ย L^* ของโดนัททอบทั้ง 4 สิ่งทดลองไม่แตกต่างกัน
 H_1 : ค่าเฉลี่ย L^* ของโดนัททอบอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

สมมติฐานข้อที่ 2 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย a^*

- H_0 : ค่าเฉลี่ย a^* ของโดนัททอบทั้ง 4 สิ่งทดลองไม่แตกต่างกัน
 H_1 : ค่าเฉลี่ย a^* ของโดนัททอบอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

สมมติฐานข้อที่ 3 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย b^*

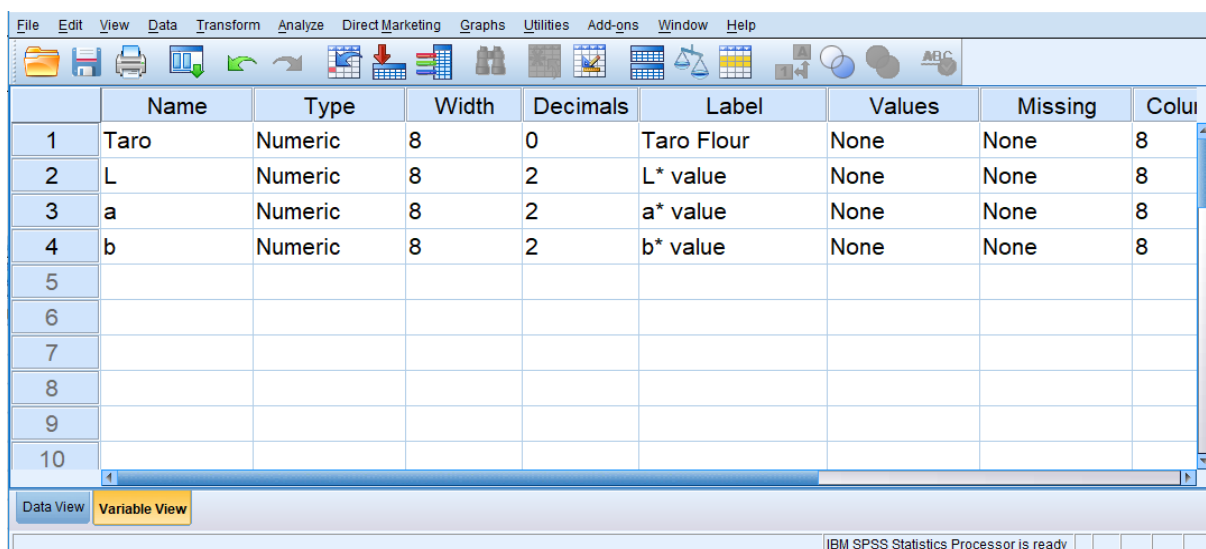
- H_0 : ค่าเฉลี่ย b^* ของโดนัททอบทั้ง 4 สิ่งทดลองไม่แตกต่างกัน
 H_1 : ค่าเฉลี่ย b^* ของโดนัททอบอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

7.7.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ One-Way ANOVA

(1) เปิดโปรแกรม SPSS ในกรณีนี้การป้อนข้อมูล ปริมาณแป้งเฟือก, ค่า L^* , ค่า a^* และค่า b^* ของโดนัทแต่ละสิ่งทดลอง ให้กำหนดรายละเอียดตัวแปรในหน้าต่าง Variable View ดังภาพที่ 7.10 ใช้ตัวแปร var จำนวน 4 ตัวแปรสำหรับป้อนข้อมูล 4 ตัว ได้แก่ ปริมาณแป้งเฟือก, ค่า L^* , ค่า a^* และ ค่า b^* ดังนั้นกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ปริมาณแป้งเฟือก	Taro	Taro Flour	0	ไม่ต้องกำหนดเพราะใช้ข้อมูลจริงในการป้อน
ค่า L^*	L	L^* value	2	
ค่า a^*	a	a^* value	2	
ค่า b^*	b	b^* value	2	

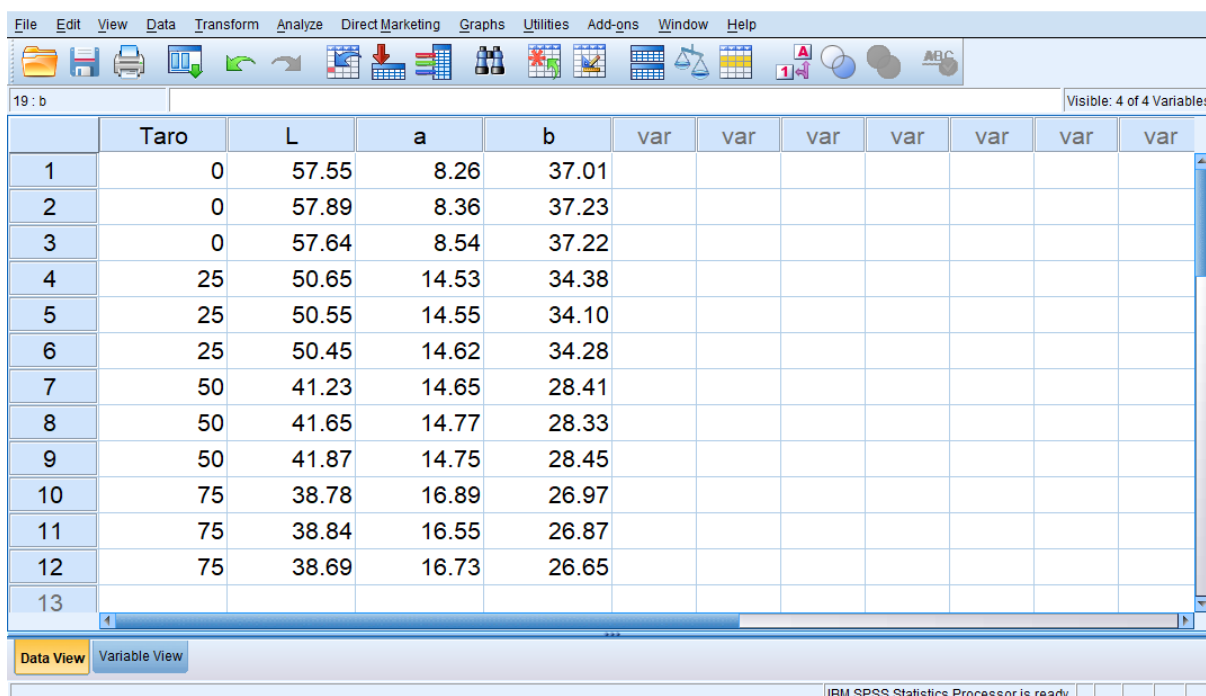
หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



	Name	Type	Width	Decimals	Label	Values	Missing	Column
1	Taro	Numeric	8	0	Taro Flour	None	None	8
2	L	Numeric	8	2	L* value	None	None	8
3	a	Numeric	8	2	a* value	None	None	8
4	b	Numeric	8	2	b* value	None	None	8
5								
6								
7								
8								
9								
10								

ภาพที่ 7.10 กำหนดรายละเอียดตัวแปรในหน้าต่าง Variable View

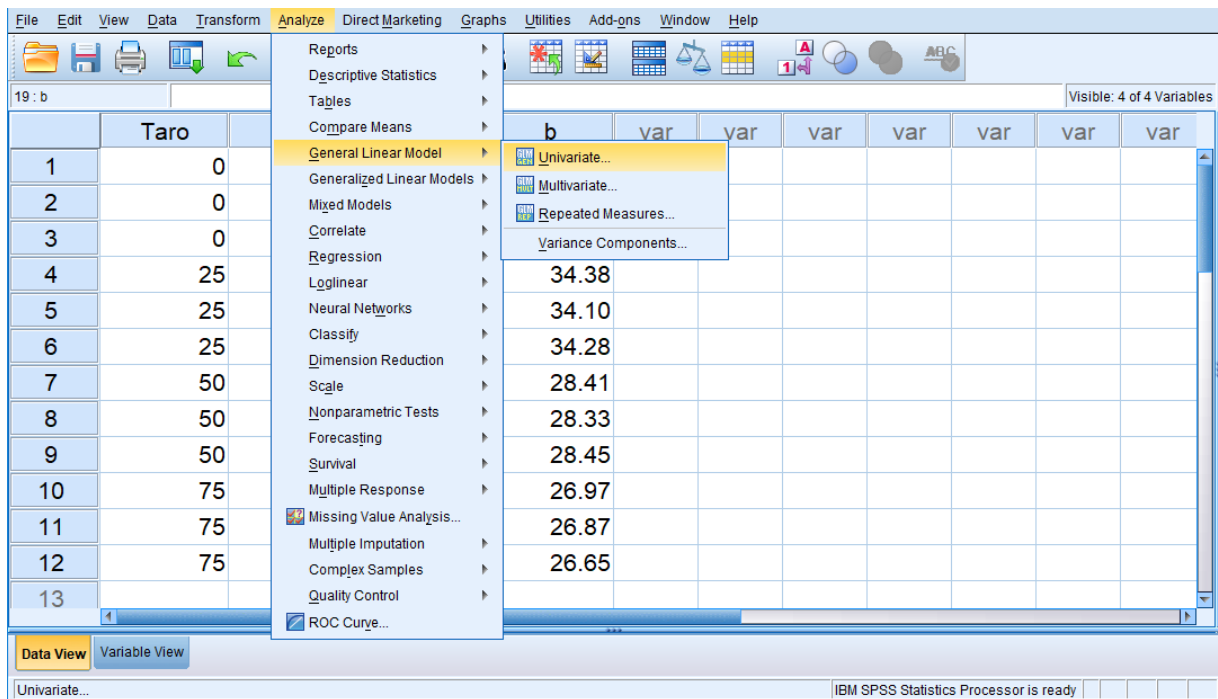
(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 7.11 โดยป้อนข้อมูล ค่าสี L*, ค่าสี a* และ ค่าสี b* เรียงตามปริมาณแป้งเผือกที่ใช้



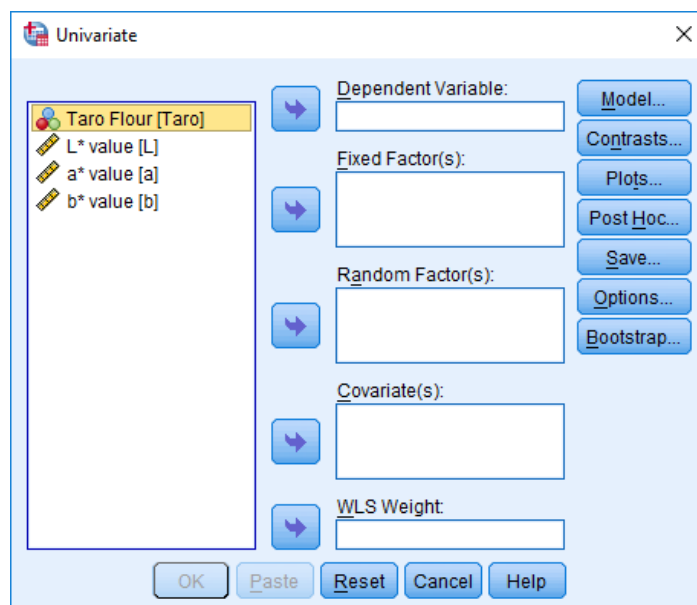
	Taro	L	a	b	var	var	var	var	var	var	var
1	0	57.55	8.26	37.01							
2	0	57.89	8.36	37.23							
3	0	57.64	8.54	37.22							
4	25	50.65	14.53	34.38							
5	25	50.55	14.55	34.10							
6	25	50.45	14.62	34.28							
7	50	41.23	14.65	28.41							
8	50	41.65	14.77	28.33							
9	50	41.87	14.75	28.45							
10	75	38.78	16.89	26.97							
11	75	38.84	16.55	26.87							
12	75	38.69	16.73	26.65							
13											

ภาพที่ 7.11 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze ➤ General Linear Model ➤ Univariate...** ดังภาพที่ 7.12 จะปรากฏหน้าจอ ดังภาพที่ 7.13 จะเห็นว่าในกล่องซ้ายมือแสดง ชื่อ Label [ชื่อ Name] ของตัวแปรแต่ละตัว

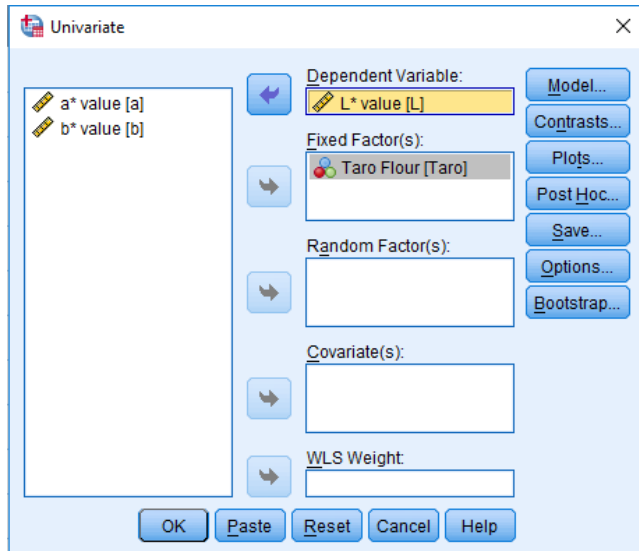


ภาพที่ 7.12 เมนูหน้าจอคำสั่ง Analyze ➤ General Linear Model ➤ Univariate...



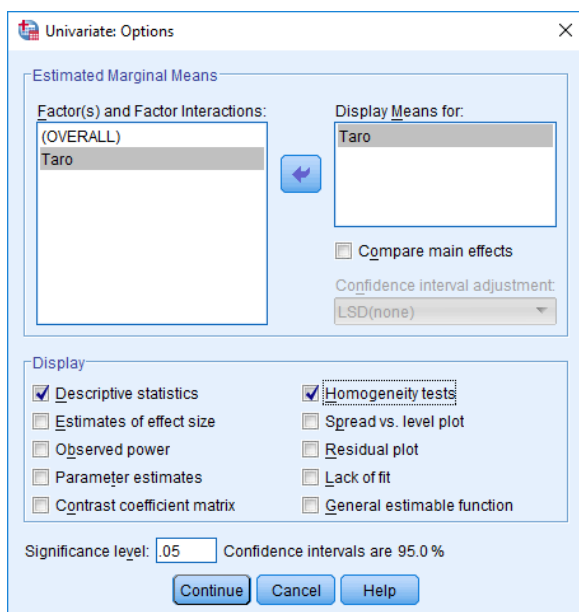
ภาพที่ 7.13 หน้าจอคำสั่ง Univariate

(4) เลือกตัวแปร Taro Flour [Taro] (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** และเลือก ตัวแปร L*value [L] (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** จะได้ดังภาพที่ 7.14 (ในกรณีมีตัวแปรตาม 3 ตัว จะต้องวิเคราะห์ความแปรปรวน 3 ครั้ง ครั้งละ 1 ตัวแปรตาม)



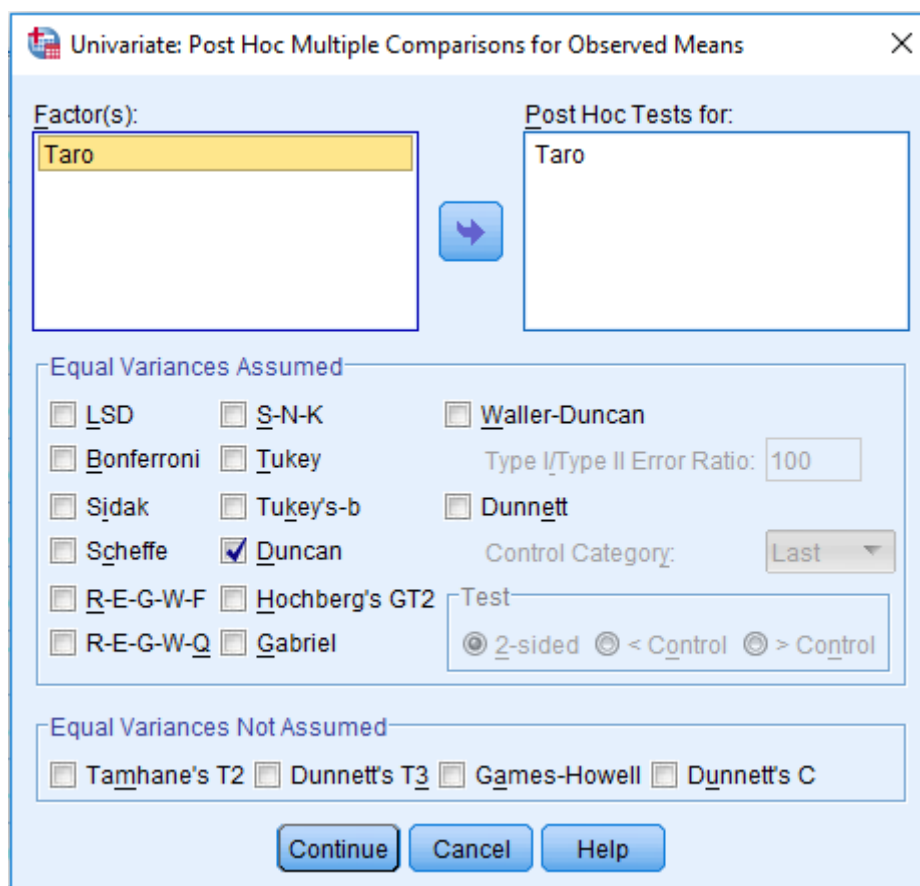
ภาพที่ 7.14 เลือกตัวแปร L*value [L] ใส่ในกล่อง Dependent Variable และ เลือกตัวแปร Taro Flour [Taro] ใส่ในกล่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Options...** จะได้หน้าจอ ดังภาพที่ 7.15 ให้กดเลือก **ตัวแปร Taro** จากกล่องซ้ายมือไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลโดนัทแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานข้อมูลโดนัทแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 7.15 หน้าจอคำสั่ง Univariate : Options

(6) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 7.16 ให้กดเลือก **ตัวแปร Taro** ที่อยู่ในกล่องซ้ายมือ (Factor(s)) ใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธีอื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 7.17



ภาพที่ 7.16 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

(7) ทำการวิเคราะห์ความแปรปรวนค่า a^* และค่า b^* ตามขั้นตอนที่ (1) – (6) จะได้ผลลัพธ์แสดงดังภาพที่ 7.18 และ 7.19 ตามลำดับ

Univariate Analysis of Variance

Between-Subjects Factors

Taro Flour	N
0	3
25	3
50	3
75	3

Descriptive Statistics

Dependent Variable: L* value

Taro Flour	Mean	Std. Deviation	N
0	57.6933	.17616	3
25	50.5500	.10000	3
50	41.5833	.32517	3
75	38.7700	.07550	3
Total	47.1492	7.81669	12

Levene's Test of Equality of Error Variances^a

Dependent Variable: L* value

F	df1	df2	Sig.
2.553	3	8	.129

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Taro

Tests of Between-Subjects Effects

Dependent Variable: L* value

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	671.802 ^a	3	223.934	5874.964	.000
Intercept	26676.527	1	26676.527	699865.160	.000
Taro	671.802	3	223.934	5874.964	.000
Error	.305	8	.038		
Total	27348.634	12			
Corrected Total	672.107	11			

a. R Squared = 1.000 (Adjusted R Squared = .999)

Estimated Marginal Means

Taro Flour

Dependent Variable: L* value

Taro Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	57.693	.113	57.433	57.953
25	50.550	.113	50.290	50.810
50	41.583	.113	41.323	41.843
75	38.770	.113	38.510	39.030

Post Hoc Tests

Taro Flour

Homogeneous Subsets

L* value

Duncan^{a, b}

Taro Flour	N	Subset			
		1	2	3	4
75	3	38.7700			
50	3		41.5833		
25	3			50.5500	
0	3				57.6933
Sig.		1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed. Based on observed means. The error term is Mean Square(Error) = .038.

a. Uses Harmonic Mean Sample Size = 3.000.

b. Alpha = .05.

ภาพที่ 7.17 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่าสี L* ด้วยคำสั่ง Univariate

จากภาพที่ 7.17 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 4 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.129 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 4 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่าเฉลี่ย L* ของทั้ง 4 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Taro มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยสี L* ของโดนัททอบอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 4 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a, b, c และ d

Univariate Analysis of Variance

Between-Subjects Factors

	N
Taro Flour 0	3
25	3
50	3
75	3

Descriptive Statistics

Dependent Variable: a* value

Taro Flour	Mean	Std. Deviation	N
0	8.3867	.14189	3
25	14.5667	.04726	3
50	14.7233	.06429	3
75	16.7233	.17010	3
Total	13.6000	3.26833	12

Levene's Test of Equality of Error Variances^a

Dependent Variable: a* value

F	df1	df2	Sig.
1.298	3	8	.340

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Taro

Tests of Between-Subjects Effects

Dependent Variable: a* value

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	117.391 ^a	3	39.130	2823.599	.000
Intercept	2219.520	1	2219.520	160157.787	.000
Taro	117.391	3	39.130	2823.599	.000
Error	.111	8	.014		
Total	2337.022	12			
Corrected Total	117.502	11			

a. R Squared = .999 (Adjusted R Squared = .999)

Estimated Marginal Means

Taro Flour

Dependent Variable: a* value

Taro Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	8.387	.068	8.230	8.543
25	14.567	.068	14.410	14.723
50	14.723	.068	14.567	14.880
75	16.723	.068	16.567	16.880

Post Hoc Tests

Taro Flour

Homogeneous Subsets

a* value

Duncan^{a, b}

Taro Flour	N	Subset		
		1	2	3
0	3	8.3867		
25	3		14.5667	
50	3		14.7233	
75	3			16.7233
Sig.		1.000	.142	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .014.

a. Uses Harmonic Mean Sample Size = 3.000.

b. Alpha = .05.

ภาพที่ 7.18 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่าสี a* ด้วยคำสั่ง Univariate

จากภาพที่ 7.18 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 4 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.340 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 4 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่าเฉลี่ย a* ของทั้ง 4 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Taro มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยสี a* ของโดนัททอบอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 3 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a, b และ c

Univariate Analysis of Variance

Between-Subjects Factors

	N
Taro Flour 0	3
25	3
50	3
75	3

Descriptive Statistics

Dependent Variable: b* value

Taro Flour	Mean	Std. Deviation	N
0	37.1533	.12423	3
25	34.2533	.14189	3
50	28.3967	.06110	3
75	26.8300	.16371	3
Total	31.6583	4.39807	12

Levene's Test of Equality of Error Variances^a

Dependent Variable: b* value

F	df1	df2	Sig.
1.037	3	8	.427

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Taro

Tests of Between-Subjects Effects

Dependent Variable: b* value

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	212.641 ^a	3	70.880	4289.278	.000
Intercept	12027.001	1	12027.001	727806.404	.000
Taro	212.641	3	70.880	4289.278	.000
Error	.132	8	.017		
Total	12239.774	12			
Corrected Total	212.773	11			

a. R Squared = .999 (Adjusted R Squared = .999)

Estimated Marginal Means

Taro Flour

Dependent Variable: b* value

Taro Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	37.153	.074	36.982	37.324
25	34.253	.074	34.082	34.424
50	28.397	.074	28.226	28.568
75	26.830	.074	26.659	27.001

Post Hoc Tests

Taro Flour

Homogeneous Subsets

b* value

Duncan^{a, b}

Taro Flour	N	Subset			
		1	2	3	4
75	3	26.8300			
50	3		28.3967		
25	3			34.2533	
0	3				37.1533
Sig.		1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed. Based on observed means. The error term is Mean Square(Error) = .017.

a. Uses Harmonic Mean Sample Size = 3.000.

b. Alpha = .05.

ภาพที่ 7.19 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่าสี b* ด้วยคำสั่ง Univariate

จากภาพที่ 7.19 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 4 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.427 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 4 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่าเฉลี่ย b* ของทั้ง 4 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Taro มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยสี b* ของโดนัทหอบอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 4 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a, b, c และ d

7.7.4 การนำเสนอผลลัพธ์ที่ได้จากการวิเคราะห์ One-Way ANOVA

จากผลลัพธ์การวิเคราะห์ความแปรปรวนของค่าสี L^* , a^* และ b^* ดังแสดงในภาพที่ 7.17, 7.18 และ 7.19 ตามลำดับนั้น นักวิจัยสามารถสรุปผลการวิเคราะห์ความแปรปรวนของค่าสีทั้ง 3 พารามิเตอร์ได้ในรูปแบบตารางเพียง 1 ตาราง โดยในการเขียนผลการวิเคราะห์ของแต่ละค่าพารามิเตอร์นั้นจะนำข้อมูลค่าเฉลี่ย (Mean) และค่าส่วนเบี่ยงเบนมาตรฐาน (Std. Deviation) ของแต่ละสิ่งทดลองจากข้อมูลในตาราง Descriptive Statistics และผลการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณจากข้อมูลในตาราง Post Hoc Tests มาใช้ในการเขียนผล ดังแสดงในตารางที่ 7.5

ตารางที่ 7.5 ค่าเฉลี่ยสี L^* , a^* และ b^* ของโดนัททอที่ใช้แป้งเผือกแทนที่แป้งสาลีในระดับที่ต่างกัน

ปริมาณแป้งเผือกที่ใช้แทนที่แป้งสาลี (%)	ค่าสี L^*	ค่าสี a^*	ค่าสี b^*
0	57.69 $\pm$ 0.18a	8.39 $\pm$ 0.14c	37.15 $\pm$ 0.12a
25	50.55 $\pm$ 0.10b	14.57 $\pm$ 0.05b	34.25 $\pm$ 0.14b
50	41.58 $\pm$ 0.33c	14.72 $\pm$ 0.06b	28.40 $\pm$ 0.06c
75	38.77 $\pm$ 0.08d	16.72 $\pm$ 0.17a	26.83 $\pm$ 0.16d

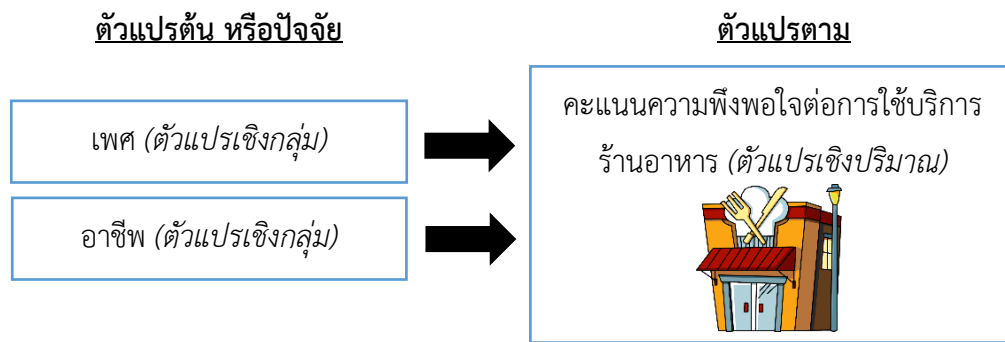
^{a-d} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

จากผลการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณในตาราง Post Hoc Tests พบว่าค่าสี L^* และค่าสี b^* แบ่งความแตกต่างออกได้เป็น 4 Subset และแต่ละ Subset มีสมาชิกในกลุ่มเพียง 1 สิ่งทดลอง ดังนั้นแต่ละสิ่งทดลองจะใช้ตัวอักษรภาษาอังกฤษเพื่อแสดงกลุ่มที่แตกต่างเพียง 1 ตัวอักษรไม่ซ้ำกัน ในขณะที่ค่าสี a^* แบ่งความแตกต่างออกได้เป็น 3 Subset (ได้แก่ a, b และ c) โดย Subset b มีสมาชิก 2 สิ่งทดลอง (ระดับแป้งเผือกที่ 25 และ 50%) ดังนั้นสองสิ่งทดลองนี้จะใช้ตัวอักษรภาษาอังกฤษเดียวกันเพื่อแสดงว่าอยู่กลุ่มเดียวกัน โดยสิ่งทดลองที่อยู่ใน Subset เดียวกัน หมายถึง ค่าเฉลี่ยไม่แตกต่างกันทางสถิติ

ข้อสังเกต ในการเรียงลำดับการใช้ตัวอักษรภาษาอังกฤษเพื่อแสดงกลุ่มที่แตกต่างของสิ่งทดลองนั้น นักวิจัยสามารถเรียงลำดับอักษร a, b, c, ... จากค่าเฉลี่ยมากไปน้อย หรือ จากค่าเฉลี่ยน้อยไปมาก ทั้งนี้ นักวิจัยต้องเลือกอย่างใดอย่างหนึ่งและต้องใช้เกณฑ์นี้กับทุกค่าตัวแปรตามในงานวิจัยเรื่องนั้นเสมอ

7.8 การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

การวิเคราะห์ความแปรปรวนแบบสองทาง หรือ Two-Way ANOVA เป็นการจำแนกข้อมูลหรือแบ่งข้อมูลด้วยตัวแปรหรือปัจจัย 2 ปัจจัย (จึงเรียกว่า Two-Way) และแต่ละปัจจัยมีหลายระดับ เช่น ผู้วิจัยคาดว่าคะแนนความพึงพอใจต่อการใช้บริการร้านอาหาร ขึ้นอยู่กับ เพศ (ชาย และ หญิง) และอาชีพ (นิสิต, อาจารย์ และ เจ้าหน้าที่) กรณีนี้จะคะแนนความพึงพอใจเป็นตัวแปรเชิงปริมาณ ส่วนเพศและอาชีพเป็นตัวแปรเชิงกลุ่ม



การวิเคราะห์ความแปรปรวนแบบสองทาง จะแบ่งความผันแปรทั้งหมดของข้อมูล ออกเป็น 3 ส่วน จากกรณีตัวอย่างดังกล่าว สามารถอธิบายได้ ดังนี้

ส่วนที่ 1 ความผันแปรของคะแนนความพึงพอใจเนื่องจากเพศต่างกัน

ส่วนที่ 2 ความผันแปรของคะแนนความพึงพอใจเนื่องจากอาชีพต่างกัน

ส่วนที่ 3 ความผันแปรของคะแนนความพึงพอใจเนื่องจากอิทธิพลร่วมของเพศและอาชีพ

โดยทั่วไปแล้วการวิเคราะห์ความแปรปรวนแบบสองทางนั้น จะมีตัวแปรตามซึ่งเป็นตัวแปรเชิงปริมาณ และตัวแปรต้นซึ่งเป็นตัวแปรเชิงกลุ่ม (หรือ ตัวแปรที่แสดงกลุ่ม) 2 ตัวแปร โดยจะเรียกตัวแปรเชิงกลุ่มว่า “ปัจจัย” เรียกปัจจัยแรกว่า “ปัจจัย A” และเรียกปัจจัยที่สองว่า “ปัจจัย B” ซึ่งถือว่าทั้งสองปัจจัยเป็น “ปัจจัยหลัก (Main factor)” และกำหนดให้ใช้อักษรภาษาอังกฤษตัวพิมพ์เล็กแสดงจำนวนระดับของปัจจัย ดังนี้

a = จำนวนระดับของปัจจัย A กรณีนี้ a = 2 (เพศ มี 2 ระดับ คือ ชาย และ หญิง)

b = จำนวนระดับของปัจจัย B กรณีนี้ b = 3 (อาชีพ มี 3 ระดับ คือ นิสิต, อาจารย์ และ เจ้าหน้าที่)

ab = จำนวนระดับของอิทธิพลร่วม (Interaction) ของปัจจัย A และ B กรณีนี้ ab = 6

การเขียนสมมติฐานทางสถิติสำหรับวิเคราะห์ความแปรปรวนแบบสองทาง

การวิเคราะห์ความแปรปรวนแบบสองทาง กรณีที่มีการวัดค่าซ้ำจะเป็นการวิเคราะห์เพื่อทดสอบผลกระทบของปัจจัยต่างๆ โดยการหาค่าตัวสถิติ F เพื่อใช้ในการทดสอบสมมติฐาน ซึ่งสามารถจำแนกได้ 3 ข้อ ดังนี้

(1) การทดสอบอิทธิพลของปัจจัย A (Main effect A) ต่อตัวแปรตาม

H_0 : ปัจจัย A ไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : ปัจจัย A มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ

H_0 : ไม่มีความแตกต่างของค่าเฉลี่ยระหว่างระดับต่างๆ ของปัจจัย A

H_1 : มีอย่างน้อย 2 ระดับของปัจจัย A ที่ค่าเฉลี่ยมีความแตกต่างกัน

(2) การทดสอบอิทธิพลของปัจจัย B (Main effect B) ต่อตัวแปรตาม

H_0 : ปัจจัย B ไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : ปัจจัย B มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ

H_0 : ไม่มีความแตกต่างของค่าเฉลี่ยระหว่างระดับต่างๆ ของปัจจัย B

H_1 : มีอย่างน้อย 2 ระดับของปัจจัย B ที่ค่าเฉลี่ยมีความแตกต่างกัน

(3) การทดสอบอิทธิพลร่วมระหว่างปัจจัย A และ B (Interaction effects AxB) ต่อตัวแปรตาม เป็นการทดสอบว่าปัจจัยทั้ง 2 ปัจจัยมีผลกระทบร่วมกันหรือไม่ เขียนสมมติฐานได้ ดังนี้

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B ต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : มีปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B ต่อค่าเฉลี่ยของตัวแปรตาม

หรือ

H_0 : ไม่มีความแตกต่างของค่าเฉลี่ยระหว่างปัจจัยร่วมของปัจจัย A และปัจจัย B

H_1 : มีความแตกต่างของค่าเฉลี่ยอย่างน้อย 2 ระดับระหว่างปัจจัยร่วมของปัจจัย A และปัจจัย B

ในการวิเคราะห์ความแปรปรวนแบบสองทางจะมีการกำหนดตัวแบบ (Model) โดยกำหนดเป็นสมการทางคณิตศาสตร์ในรูปแบบเชิงเส้นตรง ซึ่งแบ่งออกได้เป็น 3 ประเภท คือ

(1) ตัวแบบผลกระทบชนิดคงที่ (Fixed Effect Model หรือ Model I)

ตัวแบบผลกระทบชนิดคงที่ เป็นตัวแบบชนิดที่มีปัจจัยคงที่ ซึ่งเป็นการนำค่าที่เป็นไปได้ของตัวแปรต้น (หรือระดับของปัจจัย) มาวิเคราะห์ทุกค่าหรือทุกระดับของทั้ง 2 ปัจจัย ซึ่งหมายถึง ค่าของตัวแปรต้นเป็นค่าของประชากรเพราะนำมาหมดทุกค่า เช่น การทดสอบค่าเฉลี่ยเงินเดือนของอาจารย์ทุกมหาวิทยาลัย จำแนกตามตำแหน่งทางวิชาการ ตัวแปรต้น มี 2 ตัว คือ มหาวิทยาลัย และ ตำแหน่งทางวิชาการ ซึ่งทั้งสองตัวนี้ถือได้ว่าเป็นปัจจัยคงที่ทั้ง 2 ปัจจัย ไม่ว่าจะเก็บข้อมูลมาเมื่อใด มหาวิทยาลัยและตำแหน่งทางวิชาการที่นำมาวิเคราะห์จะเป็นตัวเดิมเหมือนกันทุกครั้ง จึงถือได้ว่าเป็นตัวแบบของการวิเคราะห์แบบผลกระทบชนิดคงที่

(2) ตัวแบบผลกระทบชนิดสุ่ม (Random Effect Model หรือ Model II)

ตัวแบบผลกระทบชนิดสุ่ม เป็นตัวแบบชนิดที่มีปัจจัยไม่คงที่ ซึ่งเป็นการนำค่าที่เป็นไปได้ของตัวแปรต้นหรือระดับของปัจจัยทั้ง 2 ปัจจัย มาวิเคราะห์บางค่าโดยเลือกมาอย่างสุ่ม ซึ่งหมายถึง ค่าของตัวแปรต้นเป็นค่าของตัวอย่างเพราะไม่ได้นำมาหมดทุกค่า เช่น การทดสอบค่าเฉลี่ยเงินเดือนของอาจารย์ทุกมหาวิทยาลัย จำแนกตามตำแหน่งทางวิชาการ ตัวแปรต้น มี 2 ตัว คือ มหาวิทยาลัย และ ตำแหน่งทางวิชาการ โดยเลือกมาบางมหาวิทยาลัยและบางตำแหน่งทางวิชาการแบบสุ่ม ตัวแปรต้นทั้งสองตัวนี้ถือได้ว่าเป็นปัจจัยแบบสุ่ม กล่าวคือ มหาวิทยาลัยและตำแหน่งทางวิชาการที่นำมาวิเคราะห์จะเปลี่ยนแปลงไปสำหรับการวิเคราะห์หรือเก็บข้อมูลแต่ละครั้งโดยขึ้นอยู่กับว่าในครั้งนั้นๆ สุ่มเลือกมหาวิทยาลัยและตำแหน่งทางวิชาการใด

(3) ตัวแบบผลกระทบชนิดผสม (Mixed Effect Model หรือ Model III)

ตัวแบบผลกระทบชนิดผสม เป็นตัวแบบชนิดที่มีปัจจัยหนึ่งเป็นปัจจัยคงที่และอีกปัจจัยหนึ่งเป็นปัจจัยแบบสุ่มที่มีการเลือกมาบางระดับของปัจจัย เช่น การทดสอบค่าเฉลี่ยเงินเดือนของอาจารย์ทุกมหาวิทยาลัยจำแนกตามตำแหน่งทางวิชาการ ตัวแปรต้น มี 2 ตัว คือ มหาวิทยาลัย และ ตำแหน่งทางวิชาการ ถ้าทำการเลือกมหาวิทยาลัยแบบสุ่มและเลือกตำแหน่งทางวิชาการทุกระดับ กรณีนี้มหาวิทยาลัยคือปัจจัยแบบสุ่ม และตำแหน่งทางวิชาการคือปัจจัยคงที่ นั่นคือตำแหน่งทางวิชาการไม่เปลี่ยนแปลงแต่มหาวิทยาลัยจะเปลี่ยนแปลงไปสำหรับการวิเคราะห์หรือเก็บข้อมูลแต่ละครั้ง เพราะจะขึ้นอยู่กับว่าในครั้งนั้นๆ เลือกสุ่มมหาวิทยาลัยใด

7.9 คำสั่ง SPSS ในการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

คำสั่งที่ใช้ในการวิเคราะห์ความแปรปรวนแบบสองทาง คือ

Analyze ➤ General Linear Model ➤ Univariate...

7.9.1 ตัวอย่างข้อมูล

นักวิจัยพัฒนาผลิตภัณฑ์ขนมไข่มุกรสชาติไทย โดยนักวิจัยคาดว่าคะแนนความชอบโดยรวมของผลิตภัณฑ์นี้จะขึ้นอยู่กับอาชีพและเพศของผู้ทดสอบ จึงได้ทำการประเมินความชอบผลิตภัณฑ์ขนมไข่มุกกับกลุ่มผู้ทดสอบ 3 กลุ่ม ได้แก่ นิสิต, อาจารย์ และ เจ้าหน้าที่ กลุ่มละ 16 คน แต่ละกลุ่มแบ่งเป็นเพศชาย 8 คน และเพศหญิง 8 คน โดยใช้สเกล 9-point hedonic scale (1=ชอบน้อยที่สุด และ 9=ชอบมากที่สุด) ได้ผลคะแนนความชอบโดยรวมดังแสดงในตารางที่ 7.6

ตารางที่ 7.6 คะแนนความชอบขนมไข่มุกรสชาติไทยที่ประเมินโดยผู้ทดสอบต่างอาชีพและต่างเพศกัน

เพศ	คนที่	กลุ่มนิสิต	กลุ่มอาจารย์	กลุ่มเจ้าหน้าที่
ชาย	1	5	7	6
	2	6	7	6
	3	4	7	6
	4	5	5	6
	5	4	6	7
	6	4	6	7
	7	4	6	7
	8	5	6	7
หญิง	1	5	7	7
	2	5	7	7
	3	6	7	8
	4	5	7	7
	5	5	6	8
	6	4	6	7
	7	6	6	7
	8	5	7	7

หมายเหตุ ในการประเมินความชอบทางด้านประสาทสัมผัสควรใช้ผู้ทดสอบไม่น้อยกว่า 50 คน

7.9.2 การเขียนสมมติฐาน

จากข้อมูลดังกล่าว พิจารณารายละเอียดของข้อมูลได้ ดังนี้

- (1) ตัวแปรต้นหรือปัจจัย มี 2 ปัจจัย คือ อาชีพ และ เพศ
 - ตัวแปรอาชีพ (หรือปัจจัย A) มี 3 ระดับ คือ นิสิต, อาจารย์ และ เจ้าหน้าที่
 - ตัวแปรเพศ (หรือปัจจัย B) มี 2 ระดับ คือ ชาย และ หญิง
- (2) ตัวแปรตาม มี 1 ตัว คือ ค่าคะแนนความชอบ

นักวิจัยสามารถเขียนสมมติฐานสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ ดังนี้

สมมติฐานข้อที่ 1 การทดสอบอิทธิพลของอาชีพ (Main effect A) ต่อค่าคะแนนความชอบ

- | | | |
|-------|---|--|
| H_0 | : | ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบทุกกลุ่มอาชีพไม่แตกต่างกัน |
| H_1 | : | ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบอย่างน้อย 2 กลุ่มอาชีพแตกต่างกัน |

สมมติฐานข้อที่ 2 การทดสอบอิทธิพลของเพศ (Main effect B) ต่อค่าคะแนนความชอบ

- | | | |
|-------|---|---|
| H_0 | : | ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบเพศชายและเพศหญิงไม่แตกต่างกัน |
| H_1 | : | ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบเพศชายและเพศหญิงแตกต่างกัน |

สมมติฐานข้อที่ 3 การทดสอบอิทธิพลร่วมระหว่างอาชีพและเพศ (Interaction effects AxB) ต่อค่าคะแนนความชอบ

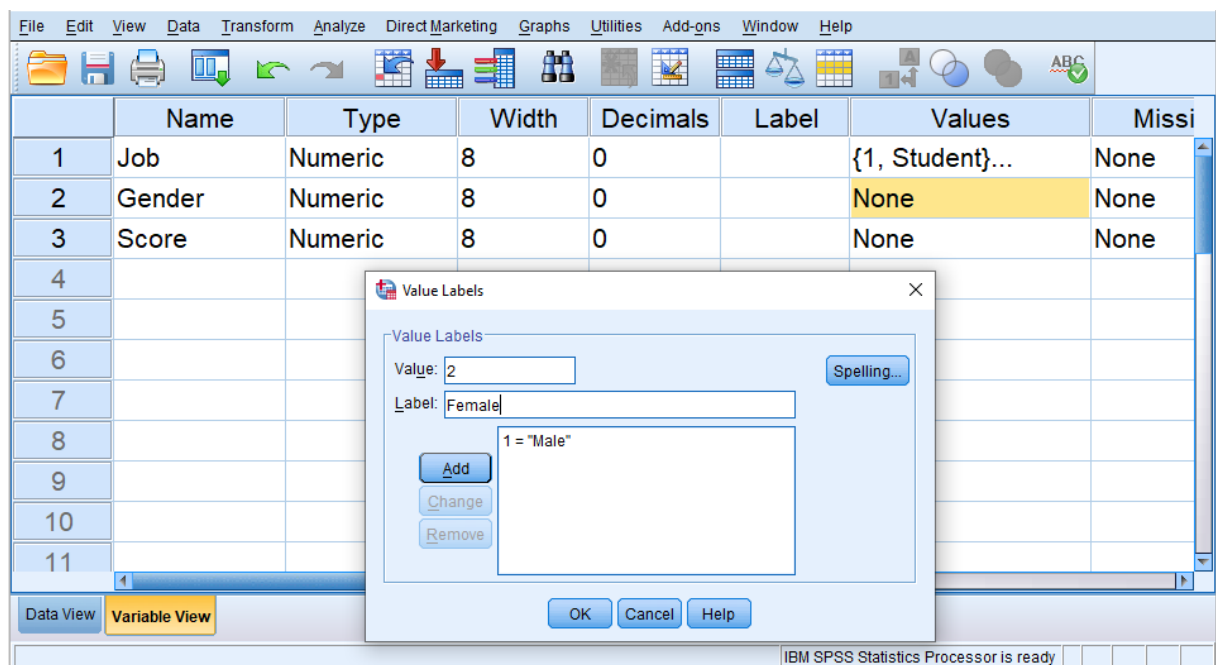
- | | | |
|-------|---|--|
| H_0 | : | ไม่มีปฏิสัมพันธ์ระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบ |
| H_1 | : | มีปฏิสัมพันธ์ระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบ |

7.9.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ความแปรปรวนแบบสองทาง

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลคะแนนความชอบของผู้ทดสอบแต่ละกลุ่ม จะต้องใช้ตัวแปร var จำนวน 3 ตัวแปรสำหรับป้อนข้อมูล 3 ตัว ได้แก่ อาชีพผู้ทดสอบ, เพศผู้ทดสอบ และ คะแนนความชอบ ดังนั้นกำหนดชื่อตัวแปรเป็น **Job**, **Gender** และ **Score** ในหน้าต่าง Variable View ดังภาพที่ 7.20 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
อาชีพผู้ทดสอบ	Job	ไม่ต้องกำหนด	0	1 = Student 2 = Teacher 3 = Staff
เพศผู้ทดสอบ	Gender	ไม่ต้องกำหนด	0	1 = Male 2 = Female
คะแนนความชอบ	Score	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด

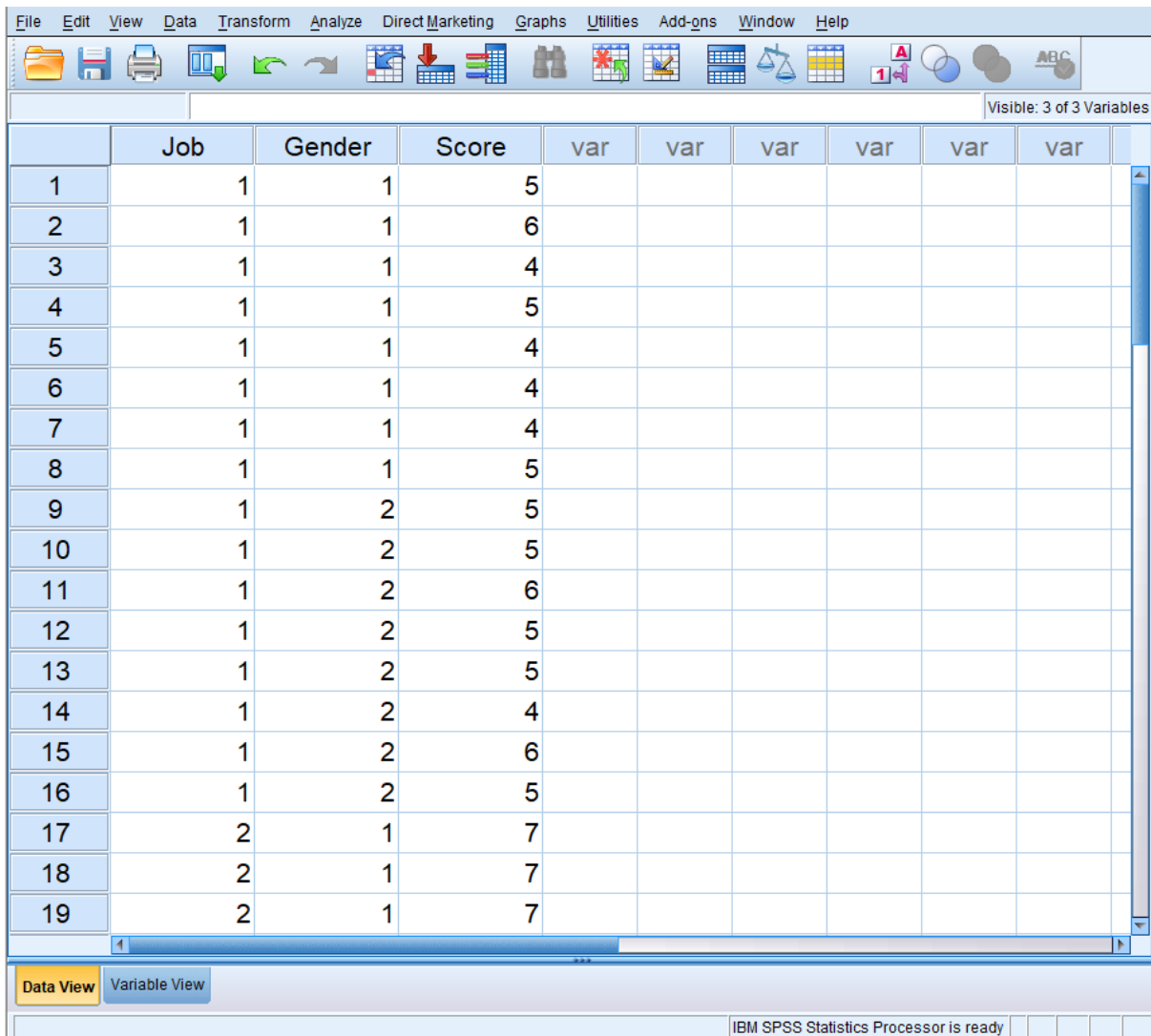
หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



ภาพที่ 7.20 กำหนดรายละเอียดตัวแปร Job, Gender และ Score ในหน้าต่าง Variable View

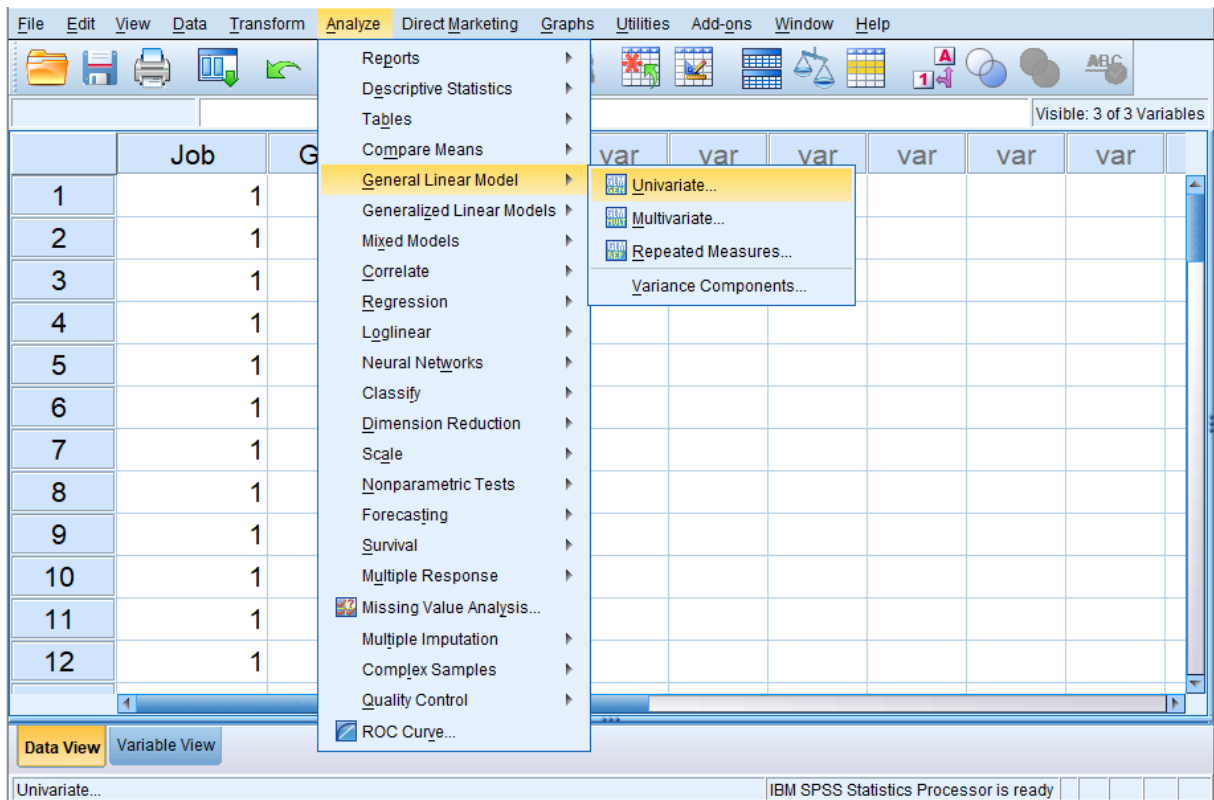
(วิธีการกำหนดค่า Values ให้ตัวแปร สามารถอ่านเพิ่มเติมได้ในบทที่ 3)

(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 7.21 โดยป้อนข้อมูลคะแนนความชอบเรียงตามกลุ่มอาชีพและเพศ

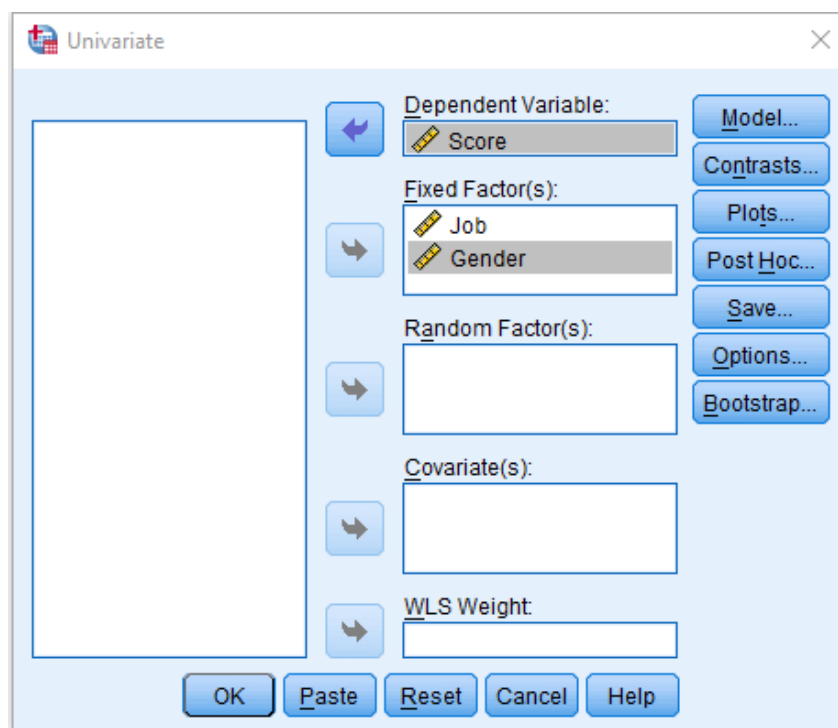


ภาพที่ 7.21 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze ➤ General Linear Model ➤ Univariate...** ดังภาพที่ 7.22 จะปรากฏหน้าจอ ดังภาพที่ 7.23

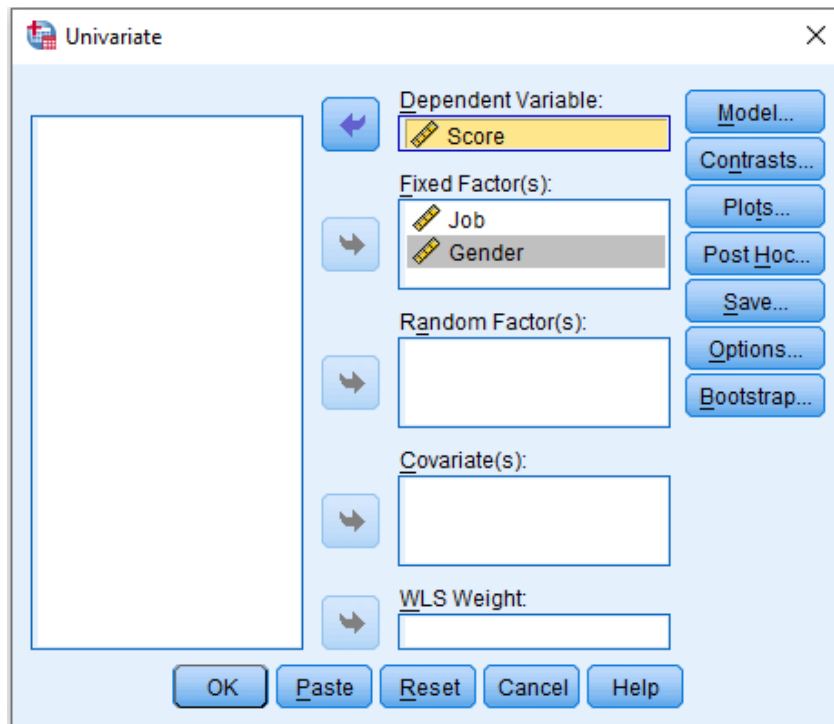


ภาพที่ 7.22 เมนูหน้าจอคำสั่ง Analyze ➤ General Linear Model ➤ Univariate...



ภาพที่ 7.23 หน้าจอคำสั่ง Univariate

(4) เลือก ตัวแปร Score (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร Job และ ตัวแปร Gender (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 7.24

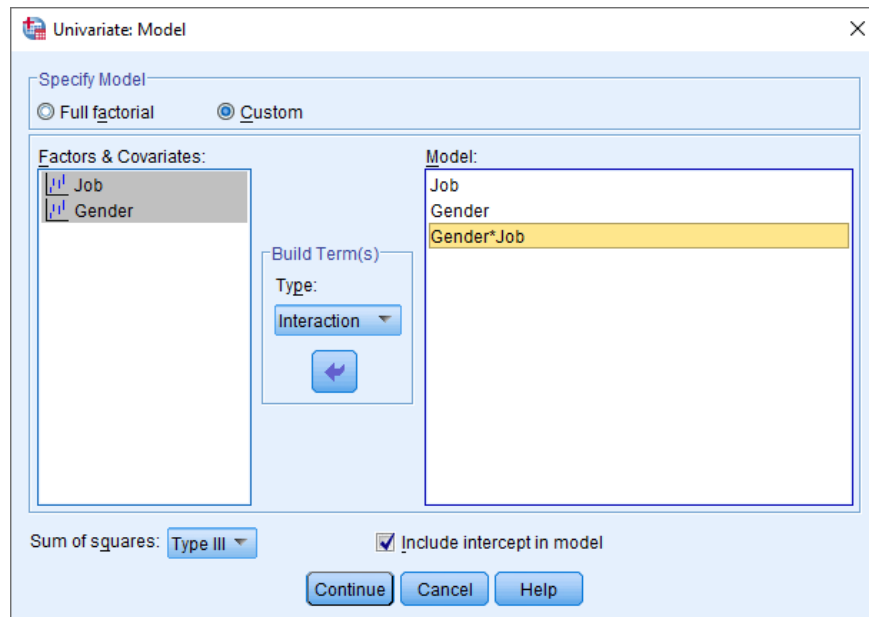


ภาพที่ 7.24 เลือกตัวแปร Score ใส่ในกล่อง Dependent Variable และ
เลือกตัวแปร Job และ Gender ใส่ในกล่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Model...** จะได้หน้าจอ ดังภาพที่ 7.25 กำหนดค่าต่างๆ ดังนี้

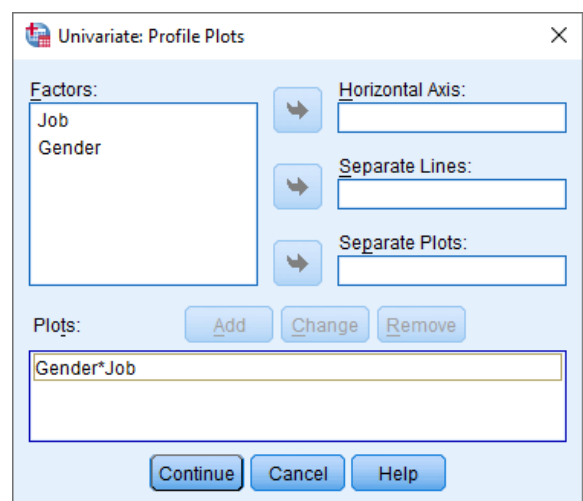
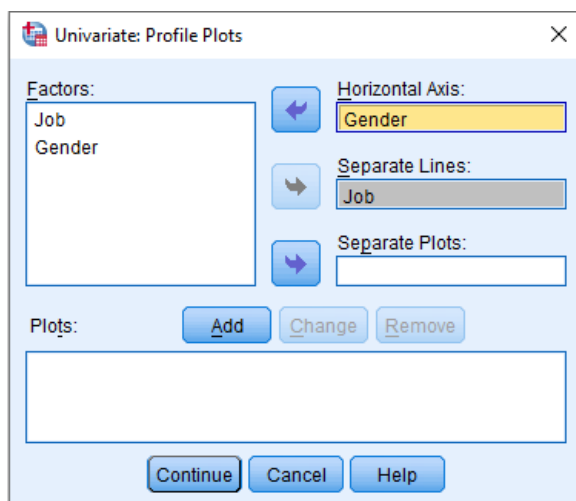
- กำหนดรูปแบบการวิเคราะห์ (Model) ที่ตำแหน่ง Specify Model ให้คลิกเลือก **Custom** ทั้งนี้รูปแบบการวิเคราะห์ (Model) มี 2 รูปแบบ Full factorial คือ รูปแบบเต็มองค์ประกอบ และ Custom คือ การกำหนดรูปแบบขึ้นเอง
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลของปัจจัยหลัก (Main Effects) โดยกด Build Term(s) เลือก **Main Effects** แล้วเลือกตัวแปร Job และ Gender ในกล่องซ้ายมือ ใส่ในกล่อง Model ที่อยู่ในกล่องขวามือ
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลร่วมระหว่างปัจจัยทั้ง 2 ปัจจัย (Interaction Effects) โดยกด Build Term(s) เลือก **Interaction** แล้วกดคีย์บอร์ด **Shift ↑** ค้างไว้ กดเลือกตัวแปร Job และ Gender (กดเลือกแบบจับคู่) ในกล่องซ้ายมือใส่ในกล่อง Model จะได้ตัวแปร **Gender*Job** ในกล่องขวามือ

- เมื่อกำหนดค่าเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



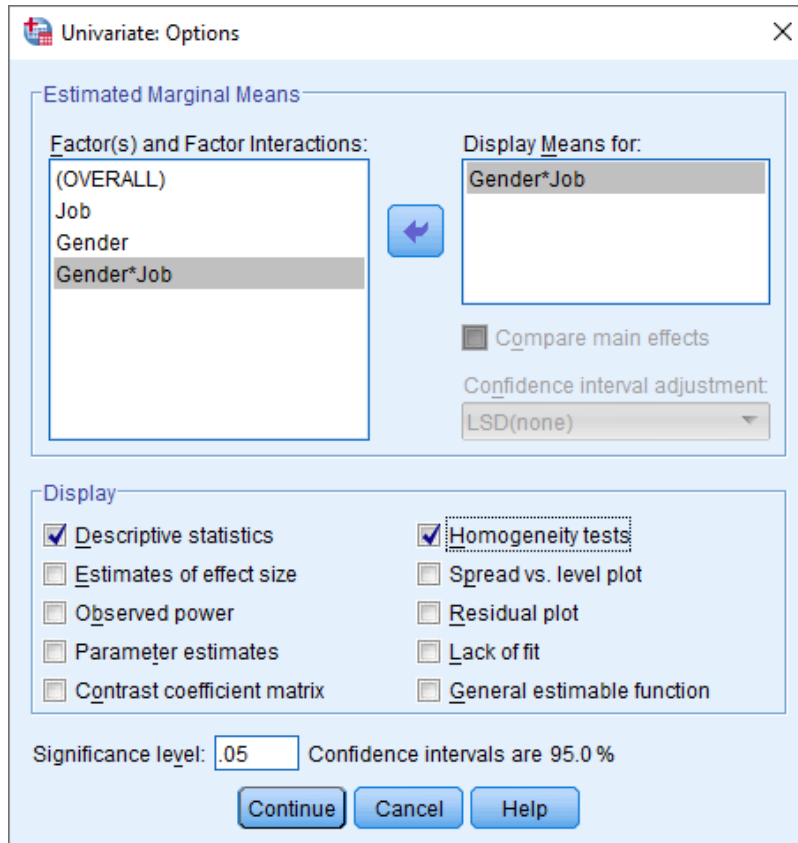
ภาพที่ 7.25 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลร่วมของทั้งสองปัจจัย จะได้หน้าจอตั้งภาพที่ 7.26 กดเลือกตัวแปร Gender ในกล่อง Horizontal Axis (แกนนอน) และตัวแปร Job ในกล่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Gender*Job** ในกล่อง Plots ด้านล่างเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 7.26 หน้าจอคำสั่ง Univariate : Profile Plots

(7) เลือกคำสั่ง **Options...** เพื่อให้แสดงค่าสถิติเบื้องต้นของแต่ละกลุ่ม จะได้หน้าจอ ดังภาพที่ 7.27 กดเลือก **ตัวแปร Gender*Job** ใส่ในกล่อง Display Means for แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) เมื่อเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate และกด OK



ภาพที่ 7.27 หน้าจอคำสั่ง Univariate : Options

(8) ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 7.28

Univariate Analysis of Variance

Between-Subjects Factors

		Value Label	N
Job	1	student	16
	2	Teacher	16
	3	Staff	16
Gender	1	Male	24
	2	Female	24

Descriptive Statistics

Dependent Variable: Score

Job	Gender	Mean	Std. Deviation	N
student	Male	4.63	.744	8
	Female	5.13	.641	8
	Total	4.88	.719	16
Teacher	Male	6.25	.707	8
	Female	6.63	.518	8
	Total	6.44	.629	16
Staff	Male	6.50	.535	8
	Female	7.25	.463	8
	Total	6.87	.619	16
Total	Male	5.79	1.062	24
	Female	6.33	1.049	24
	Total	6.06	1.080	48

Levene's Test of Equality of Error Variances^a

Dependent Variable: Score

F	df1	df2	Sig.
.754	5	42	.588

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Job + Gender + Job * Gender

Tests of Between-Subjects Effects

Dependent Variable: Score

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	39.188 ^a	5	7.838	21.067	.000
Intercept	1764.187	1	1764.187	4742.136	.000
Job	35.375	2	17.687	47.544	.000
Gender	3.521	1	3.521	9.464	.004
Job * Gender	.292	2	.146	.392	.678
Error	15.625	42	.372		
Total	1819.000	48			
Corrected Total	54.813	47			

a. R Squared = .715 (Adjusted R Squared = .681)

Estimated Marginal Means

Gender * Job

Dependent Variable: Score

Gender	Job	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
Male	student	4.625	.216	4.190	5.060
	Teacher	6.250	.216	5.815	6.685
	Staff	6.500	.216	6.065	6.935
Female	student	5.125	.216	4.690	5.560
	Teacher	6.625	.216	6.190	7.060
	Staff	7.250	.216	6.815	7.685

Profile Plots



ภาพที่ 7.28 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนแบบ Two-Way ANOVA ด้วยคำสั่ง Univariate

7.9.4 การอ่านผลลัพธ์จากการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

จากผลการวิเคราะห์ในภาพที่ 7.28 แสดงผลลัพธ์ได้ 6 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Between-Subjects Factors เป็นส่วนที่แสดงรายละเอียดของข้อมูลที่ป้อน ดังนี้

Job	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 3 กลุ่ม ได้แก่ 1, 2 และ 3
Gender	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 1 และ 2
Value Label	คือ	ค่าความหมายที่แท้จริงของแต่ละกลุ่มที่จำแนกโดยตัวแปร Job และตัวแปร Gender
N	คือ	จำนวนข้อมูลในแต่ละกลุ่มที่จำแนกโดยตัวแปร Job และตัวแปร Gender ในที่นี้ แต่ละกลุ่มที่จำแนกด้วยตัวแปร Job มีจำนวนข้อมูลเท่ากัน คือ 16 ค่า และ แต่ละกลุ่มที่จำแนกด้วยตัวแปร Gender มีจำนวนข้อมูลเท่ากัน คือ 24 ค่า

ส่วนที่ 2 ตาราง Descriptive Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบจากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Job	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 3 กลุ่ม ได้แก่ Student, Teacher, Staff
Gender	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ Male, Female
Mean	คือ	ค่าเฉลี่ยคะแนนความชอบของแต่ละกลุ่ม
Std. Deviation	คือ	ค่าส่วนเบี่ยงเบนมาตรฐานของคะแนนความชอบ ที่แสดงการกระจายของข้อมูล
N	คือ	จำนวนข้อมูลในแต่ละกลุ่มที่จำแนกด้วยตัวแปร Job และตัวแปร Gender

ส่วนที่ 3 ตาราง Levene's Test of Equality of Error Variances ใช้พิจารณาการกระจายของข้อมูลหรือความแปรปรวนของข้อมูลแต่ละกลุ่ม อธิบายได้ ดังนี้

F, df1, df2	คือ	ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในกรณีนี้การทดสอบสมมติฐานเป็นแบบสองทาง (Two-tailed test) เขียนได้ดังนี้ H_0 : ข้อมูลทุกกลุ่มมีการกระจายตัวไม่แตกต่างกัน H_1 : ข้อมูลทุกกลุ่มมีการกระจายตัวแตกต่างกัน ค่า Sig. มีค่าเท่ากับ 0.588 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทุกกลุ่มมีการกระจายตัวไม่แตกต่างกันที่ ระดับนัยสำคัญ 0.05 (Equal variances assumed)

ส่วนที่ 4 ตาราง Tests of Between-Subjects Effects เป็นส่วนที่แสดงค่าสถิติต่างๆ ของตารางวิเคราะห์ความแปรปรวน เพื่อใช้ในการทดสอบสมมติฐานทางสถิติด้วยค่าสถิติ F หรือค่าความน่าจะเป็น Sig. เพื่อใช้ทดสอบสมมติฐาน อธิบายได้ ดังนี้

Type III Sum of Squares, df, Mean Square, F คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig. คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในที่นี้จะพิจารณาค่า Sig. เพื่อทดสอบอิทธิพลของปัจจัยหลักทั้ง 2 ปัจจัย (ตัวแปร Job และตัวแปร Gender) และอิทธิพลร่วมระหว่าง 2 ปัจจัยนั้น (ตัวแปร Job*Gender) ตามสมมติฐานที่เขียนไว้

ส่วนที่ 5 ตาราง Estimated Marginal Means เป็นส่วนที่แสดงค่าเฉลี่ยโดยประมาณของตัวแปรที่นำมาทดสอบ เนื่องจากตอนวิเคราะห์เลือกตัวแปร Job*Gender ดังนั้นผลที่แสดงในตารางนี้จะเป็นค่าเฉลี่ยโดยประมาณของตัวแปร Job แต่ละกลุ่มที่แยกตามระดับตัวแปร Gender จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Mean	คือ ค่าเฉลี่ยคะแนนความชอบของแต่ละกลุ่มที่จำแนกด้วยอาชีพและเพศ
Std. Error	คือ ค่าความคลาดเคลื่อนมาตรฐานในแต่ละกลุ่ม
95% Confidence Interval	คือ ค่าที่แสดงขอบเขตบนล่างของค่าเฉลี่ยในแต่ละกลุ่มในช่วงความเชื่อมั่น 95%

ส่วนที่ 6 Profile Plots เป็นส่วนที่แสดงกราฟเส้นของตัวแปรที่ต้องการทดสอบ

ในที่นี้ Profile Plots เป็นส่วนที่แสดงกราฟเส้นของตัวแปร Score กราฟที่แสดงใช้ค่าเฉลี่ยโดยประมาณ (จากตาราง Estimated Marginal Means) มาแสดงโดยจำแนกตามตัวแปรที่ได้เลือกในขั้นตอนการวิเคราะห์ คำสั่ง Plots... ในที่นี้ได้เลือกตัวแปร Gender เป็นตัวแปรแกนนอน จึงเห็นกราฟในแกนนอนมี 2 สเกล (Male และ Female) และได้เลือกตัวแปร Job แสดงเป็นเส้นกราฟ จึงเห็นในกราฟมี 3 เส้น (Student, Teacher และ Staff) ถ้าลักษณะกราฟเป็นเส้นขนานแสดงว่า 2 ปัจจัยนั้นไม่มีอิทธิพลร่วมกัน และถ้าลักษณะเส้นกราฟเป็นแบบไม่ขนานกันแสดงว่า 2 ปัจจัยนั้นมีอิทธิพลร่วมกัน ทั้งนี้เนื่องจากการนำค่าโดยประมาณมาสร้างกราฟ การพิจารณาจากกราฟเพียงอย่างเดียวอาจเกิดความผิดพลาดได้ ในการทดสอบอิทธิพลร่วมของทั้ง 2 ปัจจัยจึงควรพิจารณาจากค่า Sig. ของตัวแปร Job*Gender ในตาราง Tests of Between-Subjects Effects และใช้กราฟ Profile Plots ในการพิจารณาแนวโน้มค่าเฉลี่ยระหว่างระดับต่างๆ ของปัจจัยทั้ง 2 ตัว

7.9.5 การสรุปผลลัพธ์การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่ได้กำหนดไว้ ดังนี้

สมมติฐานข้อที่ 1 การทดสอบอิทธิพลของอาชีพ (Main effect A) ต่อค่าคะแนนความชอบ

H_0 : ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบทุกกลุ่มอาชีพไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบอย่างน้อย 2 กลุ่มอาชีพแตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 7.28 ค่า Sig. ของตัวแปร Job มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบอย่างน้อย 2 กลุ่มอาชีพแตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 2 การทดสอบอิทธิพลของเพศ (Main effect B) ต่อค่าคะแนนความชอบ

H_0 : ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบเพศชายและเพศหญิงไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบเพศชายและเพศหญิงแตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 7.28 ค่า Sig. ของตัวแปร Gender มีค่าเท่ากับ 0.004 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบของผู้ทดสอบเพศชายและเพศหญิงแตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 3 การทดสอบอิทธิพลร่วมระหว่างอาชีพและเพศ (Interaction effects AxB) ต่อค่าคะแนนความชอบ

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบ

H_1 : มีปฏิสัมพันธ์ระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบ

จากผลลัพธ์ที่ได้ในภาพที่ 7.28 ค่า Sig. ของตัวแปร Job* Gender มีค่าเท่ากับ 0.678 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ไม่มีปฏิสัมพันธ์ระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบที่ระดับนัยสำคัญ 0.05

7.9.6 การนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

ในการเขียนนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) สามารถนำเสนอได้หลายแบบ ในที่นี้จะนำเสนอในรูปแบบตาราง 3 รูปแบบ ซึ่งนักวิจัยจะเลือกนำเสนอแบบใดนั้นขึ้นอยู่กับวัตถุประสงค์ที่ต้องการนำเสนอข้อมูล ดังนี้

แบบที่ 1 เมื่อนักวิจัยต้องการนำเสนออิทธิพลของปัจจัยหลักและอิทธิพลร่วมระหว่างปัจจัยให้นำเสนอเฉพาะค่า P-value หรือ ค่า Sig. ซึ่งแสดงอิทธิพลของปัจจัยหลัก ได้แก่ อาชีพ และ เพศ และอิทธิพลของปัจจัยร่วมระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบ โดยใช้ข้อมูลค่า Sig. จากตาราง Tests of Between-Subjects Effects ซึ่งเขียนในรูปแบบตารางได้ดังแสดงในตารางที่ 7.7

ตารางที่ 7.7 ค่า P-value แสดงอิทธิพลของอาชีพ อิทธิพลของเพศ และอิทธิพลร่วมระหว่างอาชีพและเพศ ต่อค่าเฉลี่ยคะแนนความชอบผลิตภัณฑ์ชาสมุนไพรสมุนไพรไทย

อิทธิพลของปัจจัย	ค่า P-value
อาชีพ	0.000*
เพศ	0.004*
อาชีพ x เพศ	0.678

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

แบบที่ 2 เมื่อนักวิจัยต้องการนำเสนอค่าเฉลี่ยคะแนนความชอบของแต่ละกลุ่ม และอิทธิพลของปัจจัยหลักและอิทธิพลร่วมระหว่างปัจจัยให้นำเสนอค่าเฉลี่ยและค่าส่วนเบี่ยงเบนมาตรฐานโดยใช้ข้อมูลจากตาราง Descriptive Statistics และนำเสนอค่า P-value หรือ ค่า Sig. ซึ่งแสดงอิทธิพลของปัจจัยหลัก ได้แก่ อาชีพ และ เพศ และอิทธิพลของปัจจัยร่วมระหว่างอาชีพและเพศต่อค่าเฉลี่ยคะแนนความชอบ โดยใช้ข้อมูลจากตาราง Tests of Between-Subjects Effects ซึ่งเขียนในรูปแบบตารางได้ดังแสดงในตารางที่ 7.8

ตารางที่ 7.8 ค่าเฉลี่ยคะแนนความชอบผลิตภัณฑ์ชาสมุนไพรไทยของผู้ทดสอบจำแนกตามอาชีพและเพศ

อาชีพ	เพศ	จำนวน	คะแนนความชอบ
นิสิต	ชาย	8	4.6 $\pm$ 0.7
	หญิง	8	5.1 $\pm$ 0.6
อาจารย์	ชาย	8	6.3 $\pm$ 0.7
	หญิง	8	6.6 $\pm$ 0.5
เจ้าหน้าที่	ชาย	8	6.5 $\pm$ 0.5
	หญิง	8	7.3 $\pm$ 0.5
อิทธิพลของปัจจัย			ค่า P-value
อาชีพ			0.000*
เพศ			0.004*
อาชีพ x เพศ			0.678

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

แบบที่ 3 เมื่อนักวิจัยต้องการนำเสนอค่าเฉลี่ยคะแนนความชอบของแต่ละกลุ่มพร้อมทั้งแสดงความแตกต่างของค่าเฉลี่ยแต่ละกลุ่ม และต้องการนำเสนออิทธิพลของปัจจัยหลักและอิทธิพลร่วมระหว่างปัจจัย การนำเสนอแบบที่ 3 มีความคล้ายกับแบบที่ 2 แต่เพิ่มการวิเคราะห์เปรียบเทียบความแตกต่างของค่าเฉลี่ยระหว่างกลุ่ม ซึ่งเขียนในรูปแบบตารางได้ดังแสดงในตารางที่ 7.9

ตารางที่ 7.9 ค่าเฉลี่ยคะแนนความชอบผลิตภัณฑ์ชาสมุนไพรของไทยของผู้ทดสอบจำแนกตามอาชีพและเพศ

อาชีพ	เพศ	จำนวน	คะแนนความชอบ
นิสิต	ชาย	8	$4.6 \pm 0.7c$
	หญิง	8	$5.1 \pm 0.6c$
อาจารย์	ชาย	8	$6.3 \pm 0.7b$
	หญิง	8	$6.6 \pm 0.5b$
เจ้าหน้าที่	ชาย	8	$6.5 \pm 0.5b$
	หญิง	8	$7.3 \pm 0.5a$
อิทธิพลของปัจจัย			ค่า P-value
อาชีพ			0.000*
เพศ			0.004*
อาชีพ x เพศ			0.678

^{a-c} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

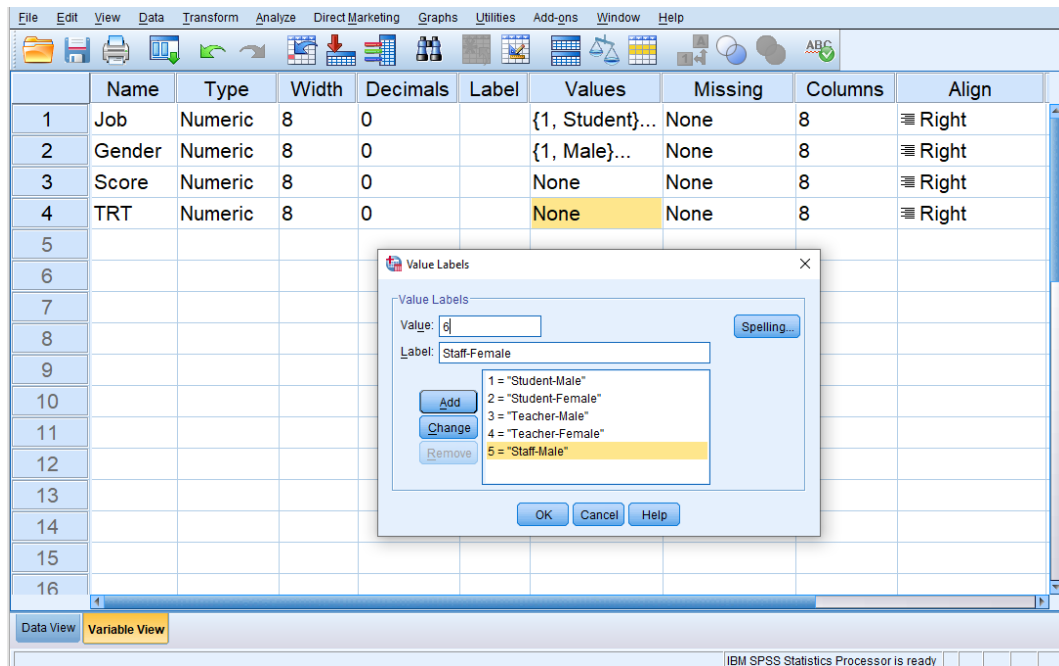
* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

การนำเสนอแบบที่ 3 ผู้วิจัยต้องวิเคราะห์ความแตกต่างของค่าเฉลี่ยระหว่างกลุ่มเพิ่มเติมทำได้ ดังนี้

(1) ผู้วิจัยต้องกำหนดตัวแปรใน Variable View ขึ้นมาใหม่อีก 1 ตัวแปร (var) กำหนดให้ชื่อ TRT ซึ่งเป็นตัวแปรที่แสดงถึงสิ่งทดลองที่เกิดจากปัจจัยร่วมระหว่างกลุ่มของอาชีพ (3 กลุ่ม) และกลุ่มของเพศ (2 กลุ่ม) หรือเรียกว่า “Treatment Combination” ในตัวอย่างนี้จะได้จำนวนสิ่งทดลอง (TRT) เท่ากับ 6 สิ่งทดลอง (คำนวณจากผลคูณของกลุ่มอาชีพและเพศ = 3×2)

(2) นำข้อมูลตัวแปร TRT (ตัวแปรต้น) และตัวแปร Score (ตัวแปรตาม) ไปวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) และเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ ด้วยคำสั่ง **Post Hoc...**

(3) กำหนดรายละเอียดและค่า Values ให้ตัวแปร TRT ในหน้าต่าง Variable View ดังภาพที่ 7.29 และป้อนข้อมูล TRT ในหน้าต่าง Data View ดังภาพที่ 7.30



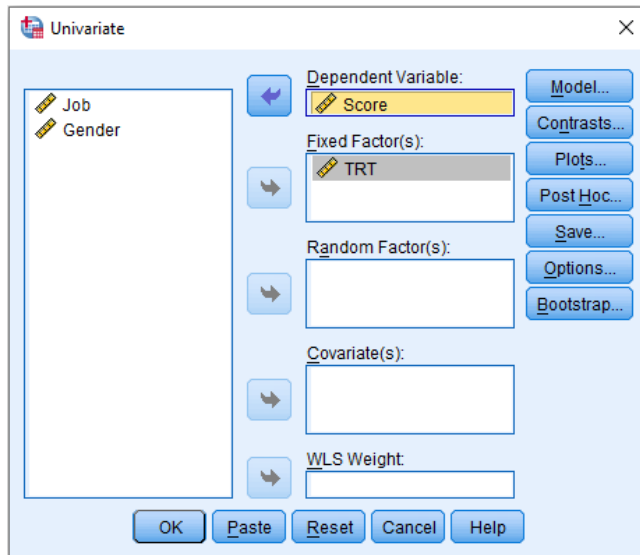
ภาพที่ 7.29 การกำหนดรายละเอียดและค่า Values ให้ตัวแปร TRT ในหน้าต่าง Variable View

(วิธีการกำหนดค่า Values ให้ตัวแปร สามารถอ่านเพิ่มเติมได้ในบทที่ 3)

File Edit View Data Transform Analyze Direct Marketing Graphs Utilities Add-ons Window Help						
Visible: 4 of 4 Variables						
	Job	Gender	Score	TRT	var	var
1	1	1	5	1		
2	1	1	6	1		
3	1	1	4	1		
4	1	1	5	1		
5	1	1	4	1		
6	1	1	4	1		
7	1	1	4	1		
8	1	1	5	1		
9	1	2	5	2		
10	1	2	5	2		
11	1	2	6	2		
12	1	2	5	2		
13	1	2	5	2		
14	1	2	4	2		
15	1	2	6	2		
16	1	2	5	2		
17	2	1	7	3		
18	2	1	7	3		

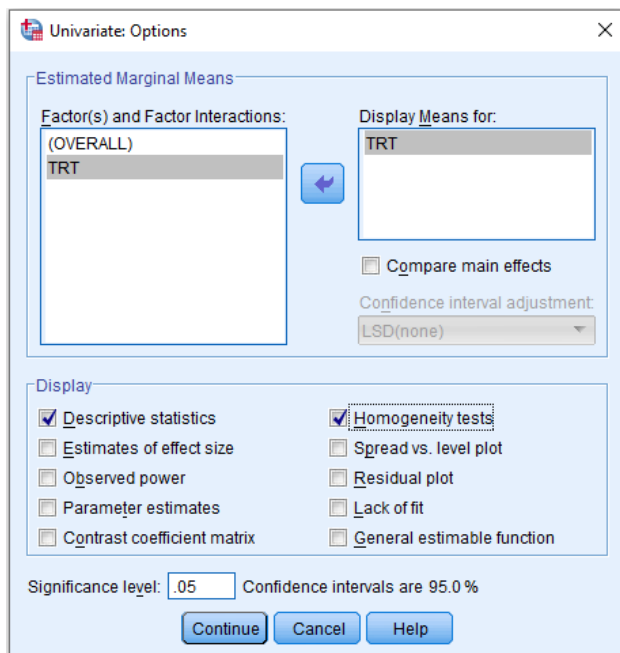
ภาพที่ 7.30 การป้อนข้อมูล TRT ในหน้าต่าง Data View

(4) เลือกเมนู **Analyze > General Linear Model > Univariate...** และเลือก ตัวแปร Score (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร TRT (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 7.31



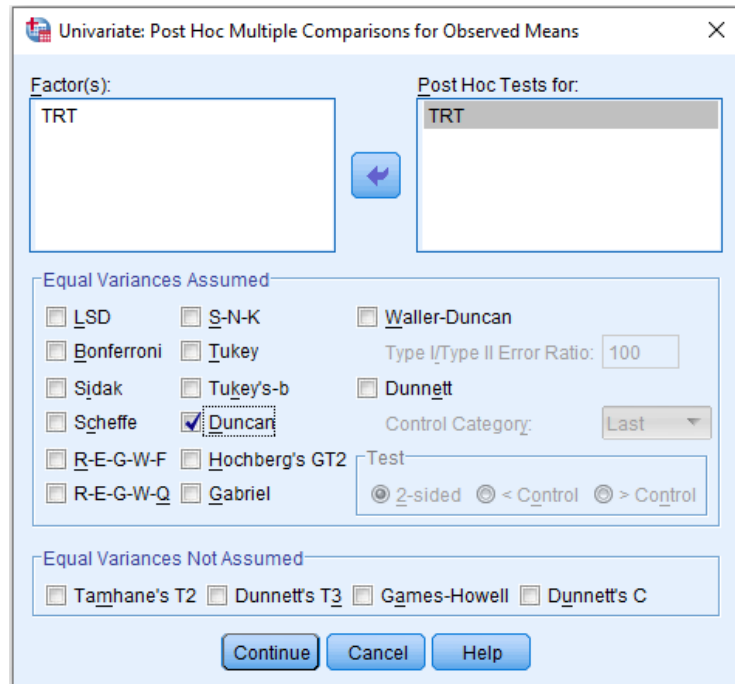
ภาพที่ 7.31 การกำหนดตัวแปร
ในเมนูคำสั่ง Univariate

(5) เลือกคำสั่ง **Options...** จะได้หน้าจอ ดังภาพที่ 7.32 กดเลือก ตัวแปร TRT ใส่ในกล่อง Display Means for ที่อยู่ด้านขวามือ และกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 7.32 การกำหนดตัวแปรในเมนู
คำสั่ง Univariate : Options

(6) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 7.33 ให้กดเลือก **ตัวแปร TRT** ในกล่อง Post Hoc Tests for: และกดเลือก **วิธี Duncan** ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 7.34



ภาพที่ 7.33 การกำหนดตัวแปรและวิธีการทดสอบในเมนูคำสั่ง Post Hoc...

Univariate Analysis of Variance

Between-Subjects Factors			
TRT		Value Label	N
1	1	Student-Male	8
	2	Student-Female	8
	3	Teacher-Male	8
	4	Teacher-Female	8
	5	Staff-Male	8
	6	Staff-Female	8

Tests of Between-Subjects Effects					
Dependent Variable: Score					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	39.188 ^a	5	7.838	21.067	.000
Intercept	1764.188	1	1764.188	4742.136	.000
TRT	39.187	5	7.837	21.067	.000
Error	15.625	42	.372		
Total	1819.000	48			
Corrected Total	54.813	47			

a. R Squared = .715 (Adjusted R Squared = .681)

Descriptive Statistics

Dependent Variable: Score

TRT	Mean	Std. Deviation	N
Student-Male	4.63	.744	8
Student-Female	5.13	.641	8
Teacher-Male	6.25	.707	8
Teacher-Female	6.63	.518	8
Staff-Male	6.50	.535	8
Staff-Female	7.25	.463	8
Total	6.06	1.080	48

Levene's Test of Equality of Error Variances^a

Dependent Variable: Score

F	df1	df2	Sig.
.754	5	42	.588

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + TRT

Estimated Marginal Means

TRT

Dependent Variable: Score

TRT	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Student-Male	4.625	.216	4.190	5.060
Student-Female	5.125	.216	4.690	5.560
Teacher-Male	6.250	.216	5.815	6.685
Teacher-Female	6.625	.216	6.190	7.060
Staff-Male	6.500	.216	6.065	6.935
Staff-Female	7.250	.216	6.815	7.685

Post Hoc Tests

TRT

Homogeneous Subsets

Score

Duncan^{a, b}

TRT	N	Subset		
		1	2	3
Student-Male	8	4.63		
Student-Female	8	5.13		
Teacher-Male	8		6.25	
Staff-Male	8		6.50	
Teacher-Female	8		6.63	
Staff-Female	8			7.25
Sig.		.109	.253	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .372.

a. Uses Harmonic Mean Sample Size = 8.000.
b. Alpha = .05.

ภาพที่ 7.34 ผลลัพธ์การวิเคราะห์ความแปรปรวนตัวแปร TRT และ ตัวแปร Score

7.10 กรณีศึกษาการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

7.10.1 ตัวอย่างข้อมูล

นิสิตปริญญาโทคนหนึ่งต้องการศึกษาผลของระดับแป้งเผือกที่ใช้แทนที่แป้งสาลีและระดับการลดปริมาณน้ำตาลที่มีผลต่อค่าความสว่าง (L^*) ของโดนัททอบ โดยศึกษาระดับของแป้งเผือกที่ใช้แทนที่แป้งสาลี 2 ระดับ คือ 25 และ 50% และ ระดับการลดปริมาณน้ำตาล 2 ระดับ คือ 0 และ 25% นำตัวอย่างโดนัททอบที่ได้ไปวัดค่าความสว่าง (L^*) ด้วยเครื่อง Chroma meter ได้ผลดังตารางที่ 7.10

ตารางที่ 7.10 ค่า L^* ของโดนัททอบเมื่อใช้แป้งเผือกแทนที่แป้งสาลีและลดปริมาณน้ำตาลที่ระดับต่างกัน

ระดับแป้งเผือกที่ใช้แทนที่ แป้งสาลีบางส่วน (%)	ระดับการลดปริมาณ น้ำตาล (%)	วัดค่าซ้ำที่	ค่าความสว่าง (L^*)
25	0	1	57.55
		2	57.89
		3	57.64
50	0	1	41.23
		2	41.65
		3	41.87
25	25	1	61.54
		2	61.35
		3	61.77
50	25	1	66.02
		2	66.81
		3	66.49

จากข้อมูลดังกล่าว พิจารณารายละเอียดได้ ดังนี้

(1) ตัวแปรต้นหรือปัจจัยที่นิสิตปริญญาโทสนใจ มี 2 ปัจจัย คือ ระดับแป้งเผือกที่ใช้ทดแทนแป้งสาลี และ ระดับการลดปริมาณน้ำตาล ดังนั้น จึงเป็นการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

- ระดับแป้งเผือกที่ใช้ทดแทนแป้งสาลี (หรือปัจจัย A) มี 2 ระดับ คือ 25 และ 50%
- ระดับการลดปริมาณน้ำตาล (หรือปัจจัย B) มี 2 ระดับ คือ 0 และ 25%

(2) ตัวแปรตาม มี 1 ตัว คือ ค่าความสว่าง L^*

7.10.2 การเขียนสมมติฐานทางสถิติสำหรับวิเคราะห์ความแปรปรวนแบบ Two-Way ANOVA

นิสิตปริญญาโทสามารถเขียนสมมติฐานสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ ดังนี้

สมมติฐานข้อที่ 1 ทดสอบอิทธิพลของระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี (Main effect A)

- H_0 : โดนนัทที่ใช้แป้งเผือกแทนที่แป้งสาลี 25 และ 50% มีค่าเฉลี่ยความสว่าง (L^*) ไม่แตกต่างกัน
 H_1 : โดนนัทที่ใช้แป้งเผือกแทนที่แป้งสาลี 25 และ 50% มีค่าเฉลี่ยความสว่าง (L^*) แตกต่างกัน

สมมติฐานข้อที่ 2 ทดสอบอิทธิพลของระดับปริมาณน้ำตาลที่ลดลง (Main effect B)

- H_0 : โดนนัทที่ลดปริมาณน้ำตาล 0 และ 25 % มีค่าเฉลี่ยความสว่าง (L^*) ไม่แตกต่างกัน
 H_1 : โดนนัทที่ลดปริมาณน้ำตาล 0 และ 25 % มีค่าเฉลี่ยความสว่าง (L^*) แตกต่างกัน

สมมติฐานข้อที่ 3 ทดสอบอิทธิพลร่วมระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลีและระดับปริมาณน้ำตาลที่ลดลง (Interaction effects AxB)

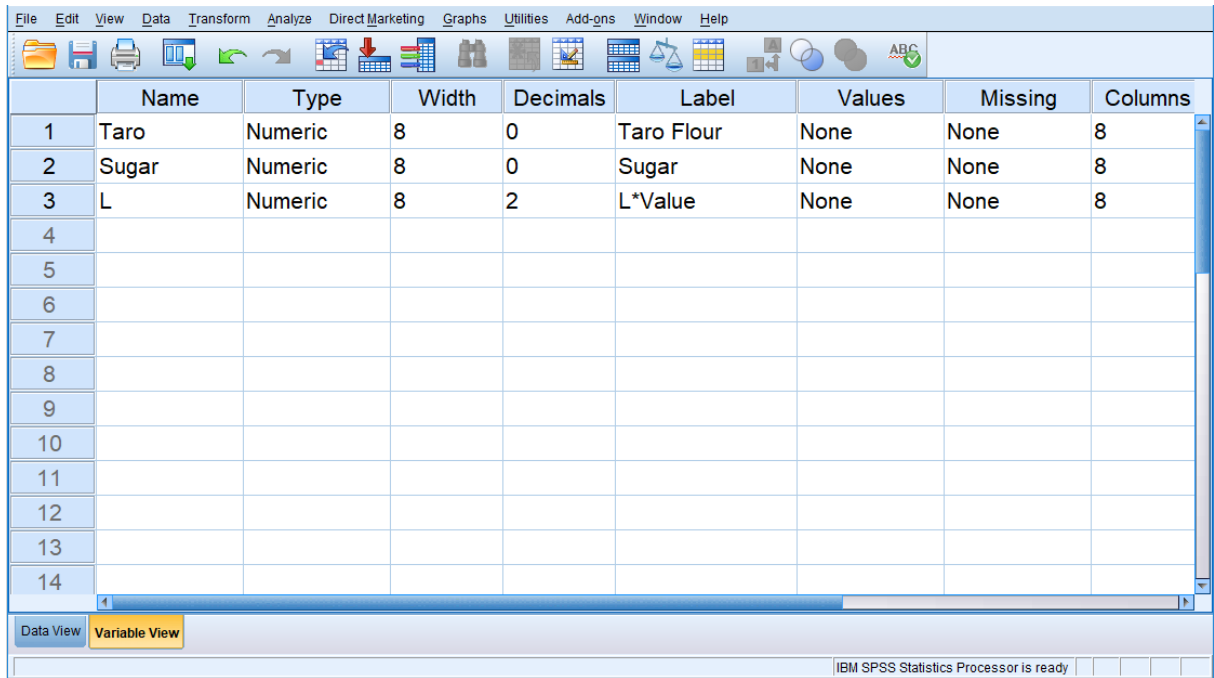
- H_0 : ไม่มีปฏิสัมพันธ์ระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี และระดับปริมาณน้ำตาลที่ลดลงต่อค่าเฉลี่ยความสว่าง (L^*) ของโดนนัท
 H_1 : มีปฏิสัมพันธ์ระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี และระดับปริมาณน้ำตาลที่ลดลงต่อค่าเฉลี่ยความสว่าง (L^*) ของโดนนัท

7.10.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ความแปรปรวนแบบ Two-Way ANOVA

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลค่าความสว่าง (L^*) ของโดนนัทแต่ละกลุ่ม จะต้องใช้ตัวแปร var จำนวน 3 ตัวแปรสำหรับป้อนข้อมูล 3 ตัว ได้แก่ ระดับแป้งเผือกที่ใช้ทดแทนแป้งสาลี, ระดับการลดปริมาณน้ำตาล และ ค่าความสว่าง (L^*) ดังนั้นกำหนดชื่อตัวแปรเป็น **Taro**, **Sugar** และ **L** ในหน้าต่าง Variable View ดังภาพที่ 7.35 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

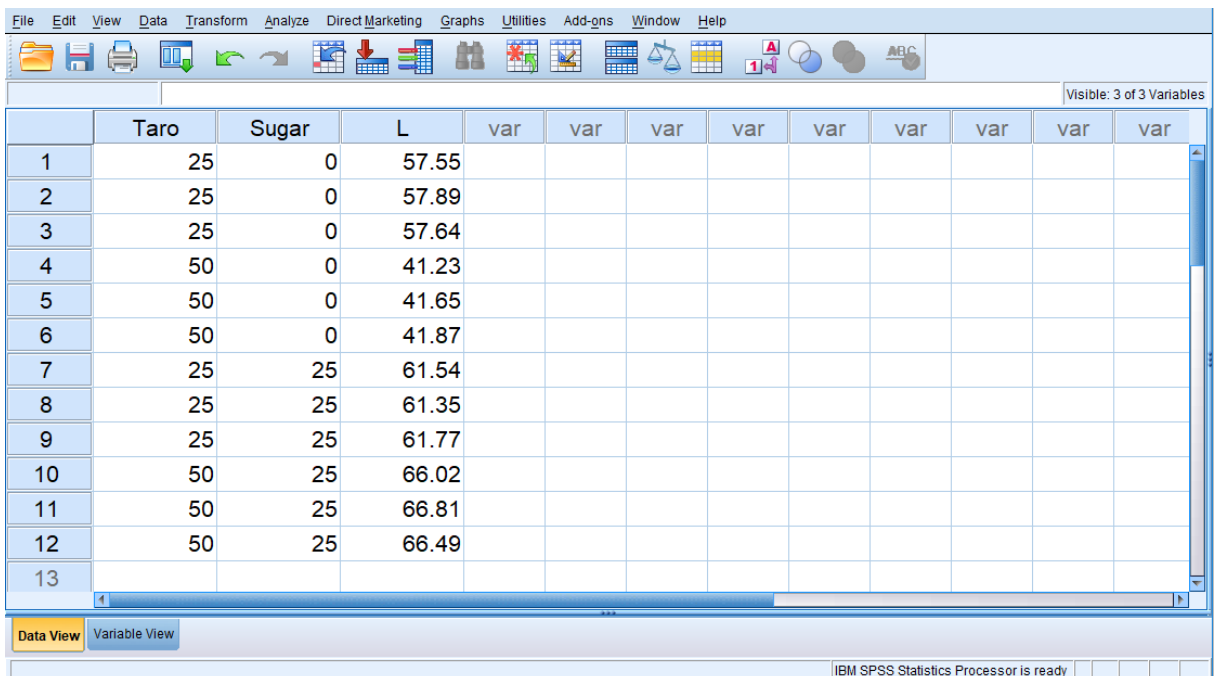
ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี	Taro	Taro Flour	0	ไม่ต้องกำหนด
ระดับการลดปริมาณน้ำตาล	Sugar	Sugar	0	ไม่ต้องกำหนด
ค่าความสว่าง (L^*)	L	L^* value	2	ไม่ต้องกำหนด

หมายเหตุ กรณีนี้กริยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



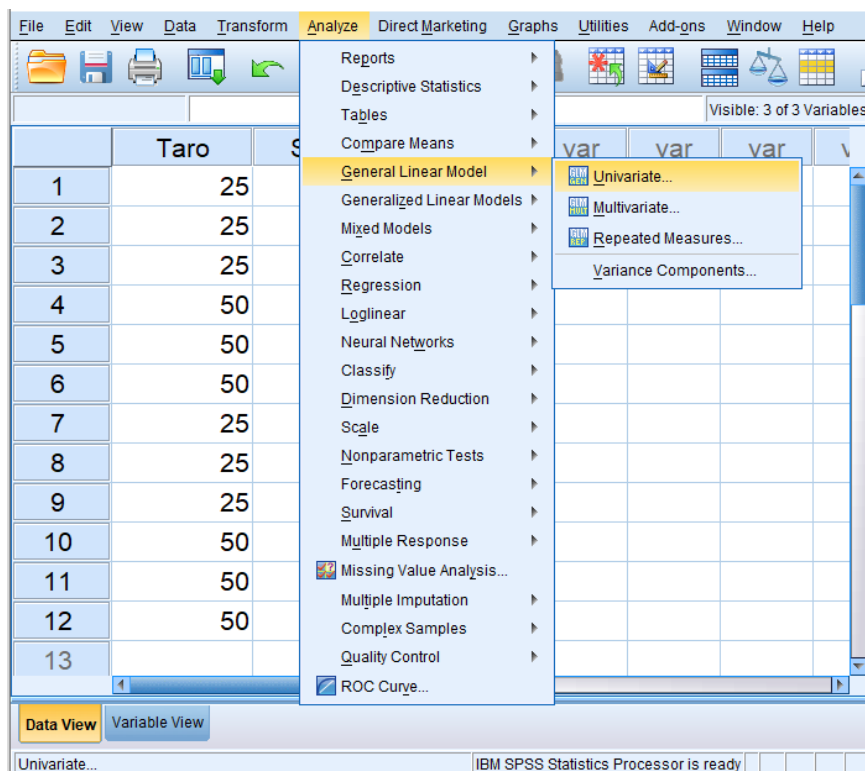
ภาพที่ 7.35 กำหนดรายละเอียดตัวแปร Taro, Sugar และ L ในหน้าต่าง Variable View

(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 7.36 โดยป้อนข้อมูลค่า L* ของโดนัทแต่ละกลุ่มเรียงตามระดับแป้งเผือกที่ใช้ทดแทนแป้งสาลี และ ระดับการลดปริมาณน้ำตาล

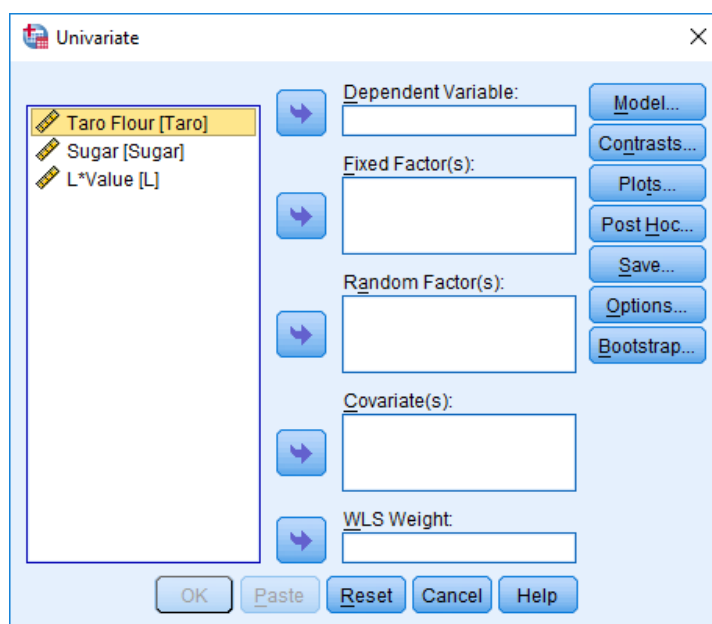


ภาพที่ 7.36 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 7.37 จะปรากฏหน้าจอตั้งภาพที่ 7.38 จะเห็นว่าในกล่องซ้ายมือแสดง ชื่อ Label [ชื่อ Name] ของตัวแปรแต่ละตัวที่กำหนดไว้

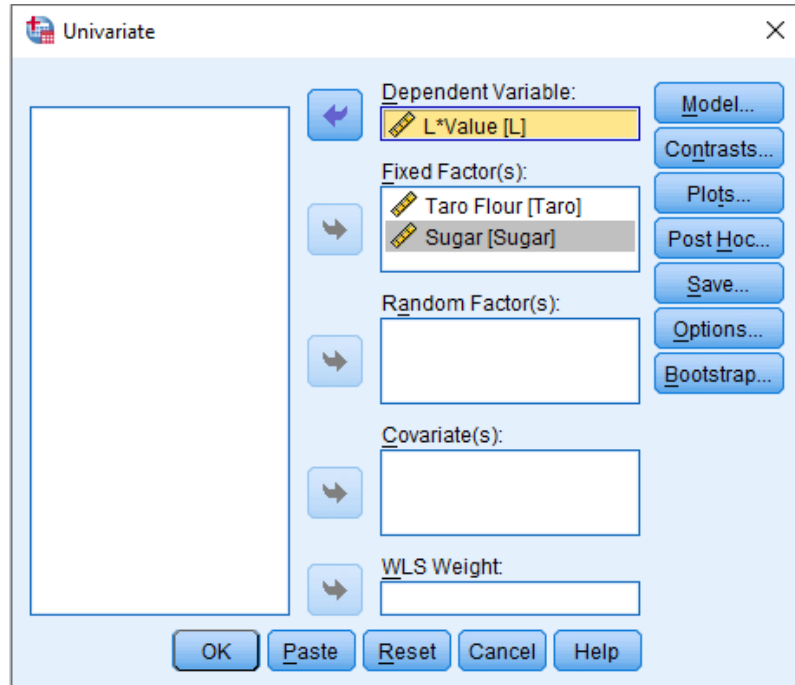


ภาพที่ 7.37 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



ภาพที่ 7.38 หน้าจอคำสั่ง Univariate

(4) เลือก ตัวแปร **L*Value [L]** (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร **Taro Flour [Taro]** และ ตัวแปร **Sugar [Sugar]** (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 7.39

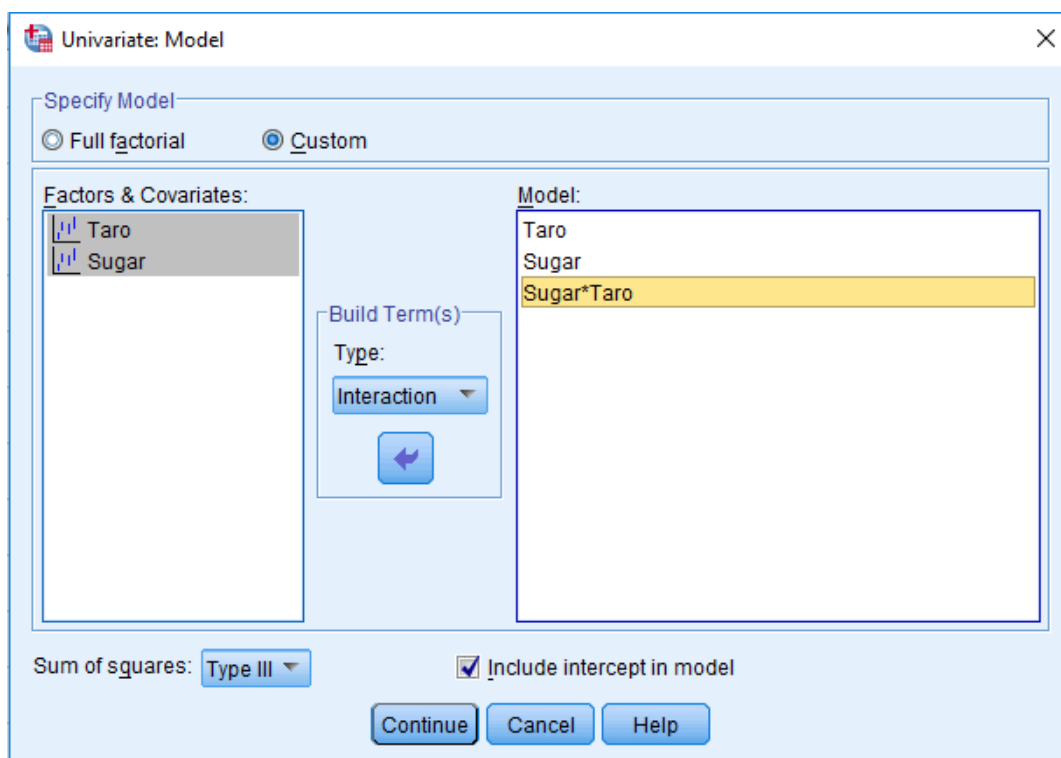


ภาพที่ 7.39 เลือกตัวแปร L*Value [L] ใส่ในกล่อง Dependent Variable และ เลือกตัวแปร Taro Flour [Taro] และ Sugar [Sugar] ใส่ในกล่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Model...** จะได้หน้าจอ ดังภาพที่ 7.40 กำหนดค่าต่างๆ ดังนี้

- กำหนดรูปแบบการวิเคราะห์ (Model) ที่ตำแหน่ง Specify Model ให้คลิกเลือก **Custom** ทั้งนี้รูปแบบการวิเคราะห์ (Model) มี 2 รูปแบบ Full factorial คือ รูปแบบเต็มองค์ประกอบ และ Custom คือ การกำหนดรูปแบบขึ้นเอง
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลของปัจจัยหลัก (Main Effects) โดยกด Build Term(s) เลือก **Main Effects** แล้วเลือกตัวแปร Taro และ ตัวแปร Sugar ในกล่องซ้ายมือใส่ในกล่อง Model ที่อยู่ในกล่องขวามือ
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลร่วมระหว่างปัจจัยทั้ง 2 ปัจจัย (Interaction Effects) โดยกด Build Term(s) เลือก **Interaction** แล้วกดคีย์บอร์ด **Shift ↑** ค้างไว้ กดเลือกตัวแปร Taro และ ตัวแปร Sugar (กดเลือกแบบจับคู่) ในกล่องซ้ายมือใส่ในกล่อง Model จะได้ตัวแปร **Sugar*Taro** ในกล่องขวามือ

- เมื่อกำหนดค่าเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



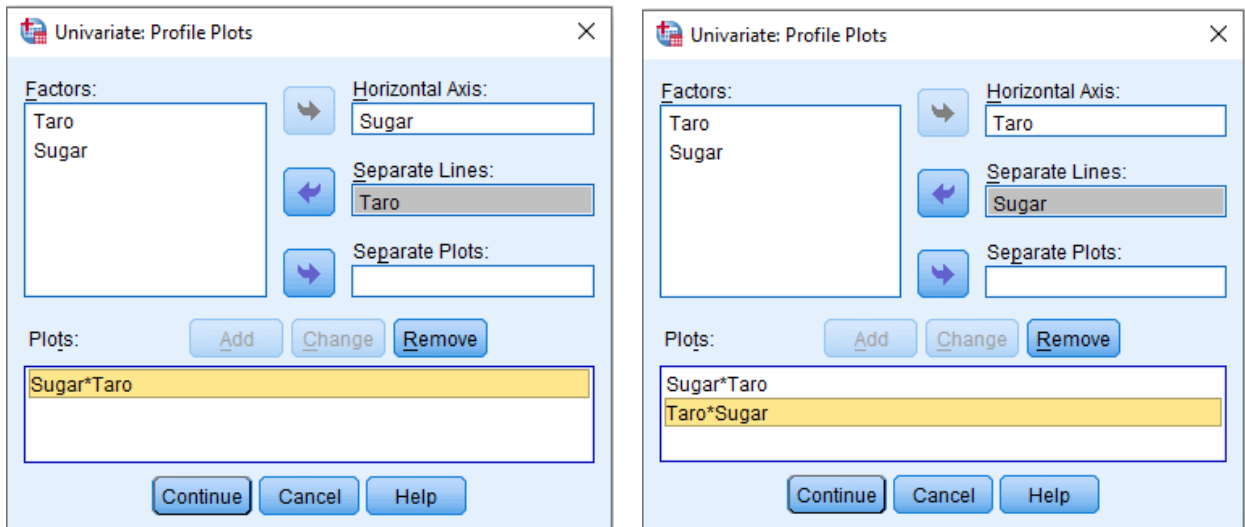
ภาพที่ 7.40 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลร่วมของทั้งสองปัจจัย จะได้หน้าจอตั้งภาพที่ 7.41 ในตัวอย่างนี้จะสร้าง 2 กราฟ โดยสลับตัวแปรที่ใช้เป็นแกนนอน และที่ใช้แสดงรูปแบบเส้นกราฟ

- กราฟที่ 1 กดเลือกตัวแปร Sugar ในกล่อง Horizontal Axis (แกนนอน) และตัวแปร Taro ในกล่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Sugar*Taro** ในกล่องด้านล่าง

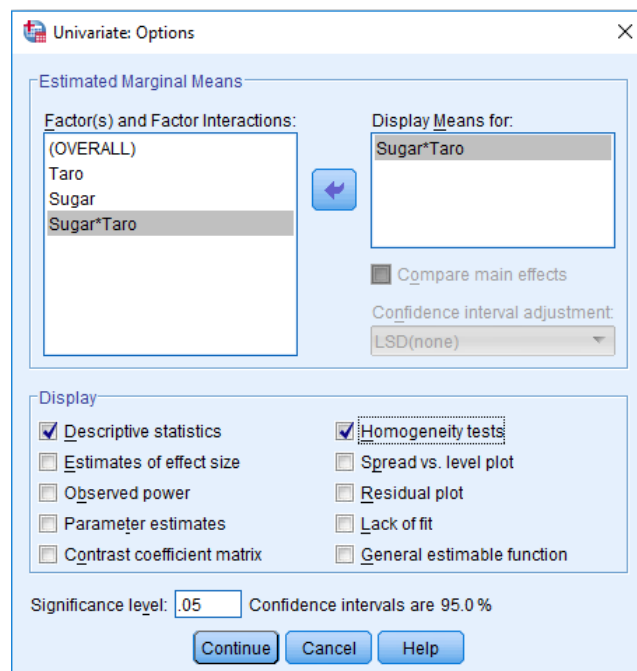
- กราฟที่ 2 กดเลือกตัวแปร Taro ในกล่อง Horizontal Axis (แกนนอน) และตัวแปร Sugar ในกล่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Taro*Sugar** ในกล่องด้านล่าง

- เมื่อกำหนดตัวแปรเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 7.41 หน้าจอคำสั่ง Profile Plots

(7) เลือกคำสั่ง **Options...** เพื่อให้แสดงค่าสถิติเบื้องต้นของแต่ละกลุ่ม จะได้หน้าจอตั้งภาพที่ 7.42 กดเลือก **ตัวแปร Sugar*Taro** ใส่ในกล่อง Display Means for แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) เมื่อเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate และกด OK จะได้ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 7.43



ภาพที่ 7.42 หน้าจอคำสั่ง Univariate: Options

Univariate Analysis of Variance

Between-Subjects Factors

	N
Taro Flour	25
50	6
Sugar	0
25	6

Descriptive Statistics

Dependent Variable: L*Value

Taro Flour	Sugar	Mean	Std. Deviation	N
25	0	57.6933	.17616	3
25	25	61.5533	.21032	3
Total		59.6233	2.12132	6
50	0	41.5833	.32517	3
50	25	66.4400	.39737	3
Total		54.0117	13.61843	6
Total	0	49.6383	8.82691	6
Total	25	63.9967	2.69160	6
Total		56.8175	9.74344	12

Levene's Test of Equality of Error Variances^a

Dependent Variable: L*Value

F	df1	df2	Sig.
.764	3	8	.545

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Taro + Sugar + Taro * Sugar

Tests of Between-Subjects Effects

Dependent Variable: L*Value

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	1043.603 ^a	3	347.868	4105.843	.000
Intercept	38738.740	1	38738.740	457229.149	.000
Taro	94.472	1	94.472	1115.048	.000
Sugar	618.485	1	618.485	7299.914	.000
Taro * Sugar	330.645	1	330.645	3902.567	.000
Error	.678	8	.085		
Total	39783.020	12			
Corrected Total	1044.280	11			

a. R Squared = .999 (Adjusted R Squared = .999)

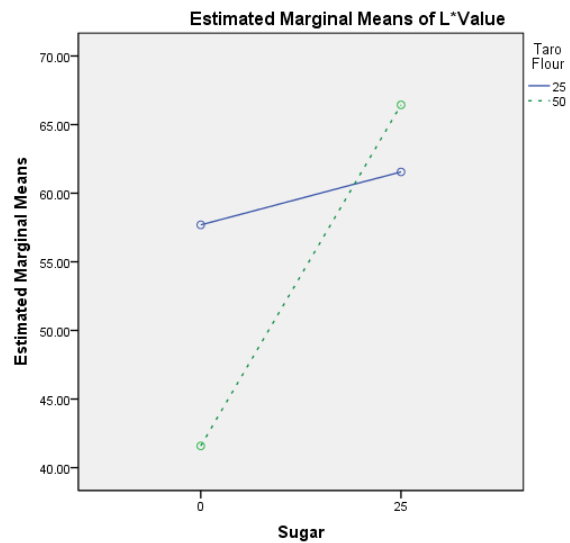
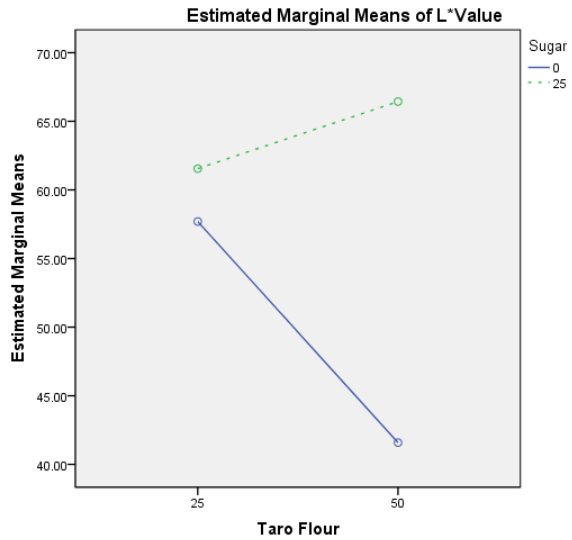
Estimated Marginal Means

Sugar * Taro Flour

Dependent Variable: L*Value

Sugar	Taro Flour	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
0	25	57.693	.168	57.306	58.081
0	50	41.583	.168	41.196	41.971
25	25	61.553	.168	61.166	61.941
25	50	66.440	.168	66.052	66.828

Profile Plots



ภาพที่ 7.43 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนแบบ Two-Way ANOVA ด้วยคำสั่ง Univariate

7.10.4 การสรุปผลลัพธ์การวิเคราะห์ความแปรปรวนแบบ Two-Way ANOVA

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่ได้กำหนดไว้ ดังนี้

สมมติฐานข้อที่ 1 ทดสอบอิทธิพลของระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี (Main effect A)

H_0 : โดนต์ที่ใช้แป้งเผือกแทนที่แป้งสาลี 25 และ 50% มีค่าเฉลี่ยความสว่าง (L^*) ไม่แตกต่างกัน

H_1 : โดนต์ที่ใช้แป้งเผือกแทนที่แป้งสาลี 25 และ 50% มีค่าเฉลี่ยความสว่าง (L^*) แตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 7.43 ค่า Sig. ของตัวแปร Taro มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ โดนต์ที่ใช้แป้งเผือกแทนที่แป้งสาลี 25 และ 50% มีค่าเฉลี่ยความสว่าง (L^*) แตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 2 ทดสอบอิทธิพลของระดับปริมาณน้ำตาลที่ลดลง (Main effect B)

H_0 : โดนต์ที่ลดปริมาณน้ำตาล 0 และ 25 % มีค่าเฉลี่ยความสว่าง (L^*) ไม่แตกต่างกัน

H_1 : โดนต์ที่ลดปริมาณน้ำตาล 0 และ 25 % มีค่าเฉลี่ยความสว่าง (L^*) แตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 7.43 ค่า Sig. ของตัวแปร Sugar มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ โดนต์ที่ลดปริมาณน้ำตาล 0 และ 25 % มีค่าเฉลี่ยความสว่าง (L^*) แตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 3 ทดสอบอิทธิพลร่วมระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลีและระดับปริมาณน้ำตาลที่ลดลง (Interaction effects AxB)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี และระดับปริมาณน้ำตาลที่ลดลงต่อค่าเฉลี่ยความสว่าง (L^*) ของโดนต์

H_1 : มีปฏิสัมพันธ์ระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี และระดับปริมาณน้ำตาลที่ลดลงต่อค่าเฉลี่ยความสว่าง (L^*) ของโดนต์

จากผลลัพธ์ที่ได้ในภาพที่ 7.43 ค่า Sig. ของตัวแปร Taro*Sugar มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ มีปฏิสัมพันธ์ระหว่างระดับแป้งเผือกที่ใช้แทนที่แป้งสาลี และระดับปริมาณน้ำตาลที่ลดลงต่อค่าเฉลี่ยความสว่าง (L^*) ของโดนต์ ที่ระดับนัยสำคัญ 0.05 และเมื่อพิจารณาจากกราฟ Profile Plots ลักษณะเส้นกราฟมีรูปแบบไม่ขนานแสดงให้เห็นว่ามีอิทธิพลร่วมระหว่างปัจจัย

7.10.5 การนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA)

ในการเขียนนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) สามารถนำเสนอได้หลายแบบ ในที่นี้จะนำเสนอในรูปแบบตาราง 2 รูปแบบดังตารางที่ 7.11 และ 7.12 ซึ่งนักวิจัยจะเลือกนำเสนอแบบใดนั้นขึ้นอยู่กับวัตถุประสงค์ที่ต้องการนำเสนอข้อมูล

ตารางที่ 7.11 ค่า P-value แสดงอิทธิพลของระดับแป้งเผือก อิทธิพลของระดับการลดน้ำตาล และอิทธิพลร่วมระหว่างระดับแป้งเผือกและระดับการลดน้ำตาลต่อค่าเฉลี่ยความสว่าง (L*) ของโดนัททอบ

อิทธิพลของปัจจัย	ค่า P-value
ระดับแป้งเผือก	0.000*
ระดับลดน้ำตาล	0.000*
ระดับแป้งเผือก × ระดับการลดน้ำตาล	0.000*

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 7.12 ค่าเฉลี่ยความสว่าง L* ของโดนัททอบที่ใช้ระดับของแป้งเผือกและระดับของการลดน้ำตาลต่างกัน

ระดับแป้งเผือกที่ใช้ทดแทนแป้งสาลี (%)	ระดับการลดน้ำตาล (%)	ค่าความสว่าง L*
25	0	57.69 ± 0.18c
50	0	41.58 ± 0.33d
25	25	61.55 ± 0.21b
50	25	66.44 ± 0.40a
อิทธิพลของปัจจัย		ค่า P-value
ระดับของแป้งเผือก		0.000*
ระดับการลดน้ำตาล		0.000*
ระดับของแป้งเผือก × ระดับของการลดน้ำตาล		0.000*

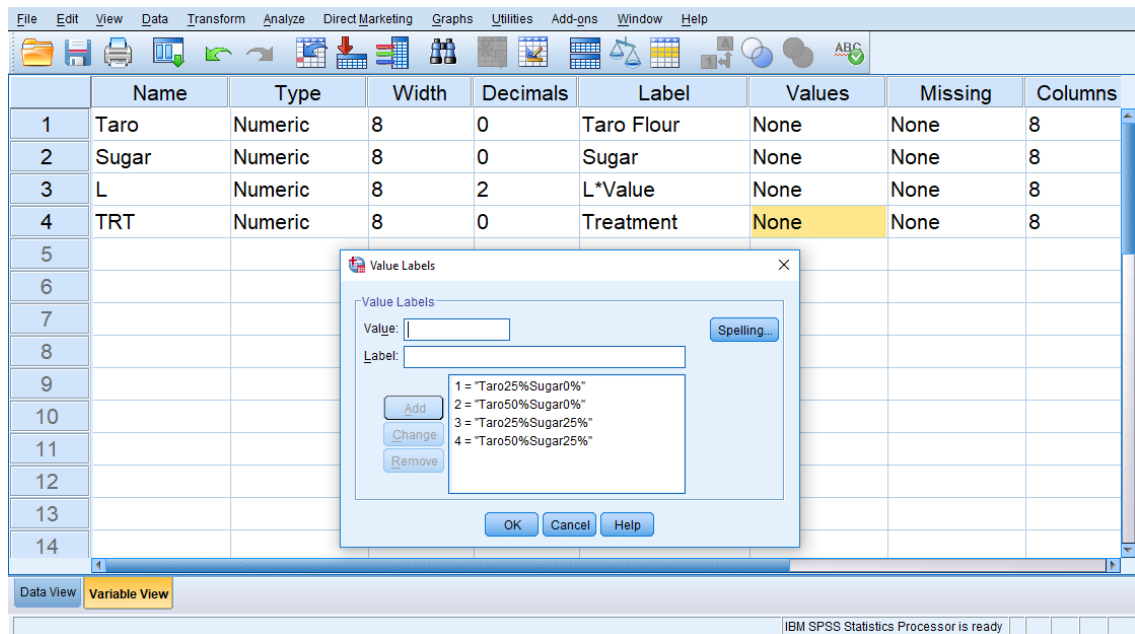
^{a-d} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 7.12 ได้มาจากการกำหนดตัวแปร var ใหม่ชื่อ **TRT** ซึ่งเป็นตัวแปรที่แสดงถึงสิ่งทดลองที่เกิดจากปัจจัยร่วมระหว่างระดับของแป้งเผือก (2 ระดับ) และระดับของการลดน้ำตาล (2 ระดับ) หรือเรียกว่า **“Treatment Combination”** ในตัวอย่างนี้จะได้จำนวนสิ่งทดลอง (TRT) เท่ากับ 4 สิ่งทดลอง (คำนวณจาก

ผลคูณของระดับแต่ละปัจจัย) นำข้อมูลตัวแปร TRT และตัวแปร L ไปวิเคราะห์ด้วยคำสั่ง **Post Hoc...** เพื่อเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ ทำเช่นเดียวกับการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) ตามขั้นตอนดังนี้

(1) กำหนดรายละเอียดตัวแปร TRT ในหน้าต่าง Variable View ดังภาพที่ 7.44 และป้อนข้อมูล TRT ในหน้าต่าง Data View ดังภาพที่ 7.45

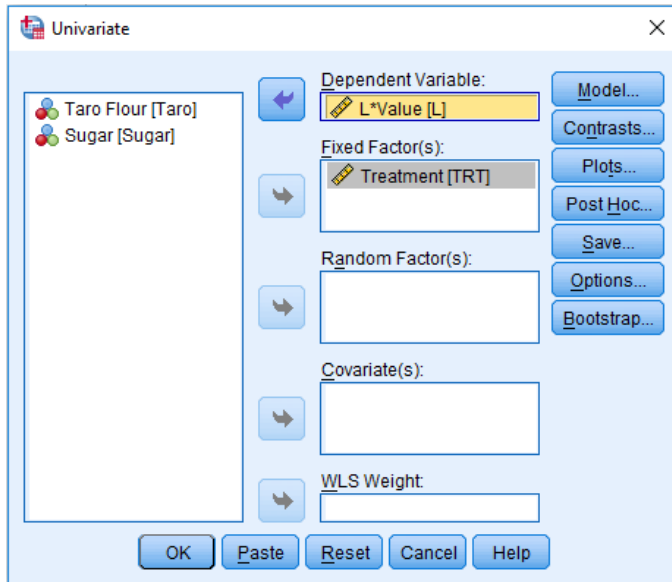


ภาพที่ 7.44 การกำหนดรายละเอียดตัวแปรใหม่ในหน้าต่าง Variable View

	Taro	Sugar	L	TRT
1	25	0	57.55	1
2	25	0	57.89	1
3	25	0	57.64	1
4	50	0	41.23	2
5	50	0	41.65	2
6	50	0	41.87	2
7	25	25	61.54	3
8	25	25	61.35	3
9	25	25	61.77	3
10	50	25	66.02	4
11	50	25	66.81	4
12	50	25	66.49	4

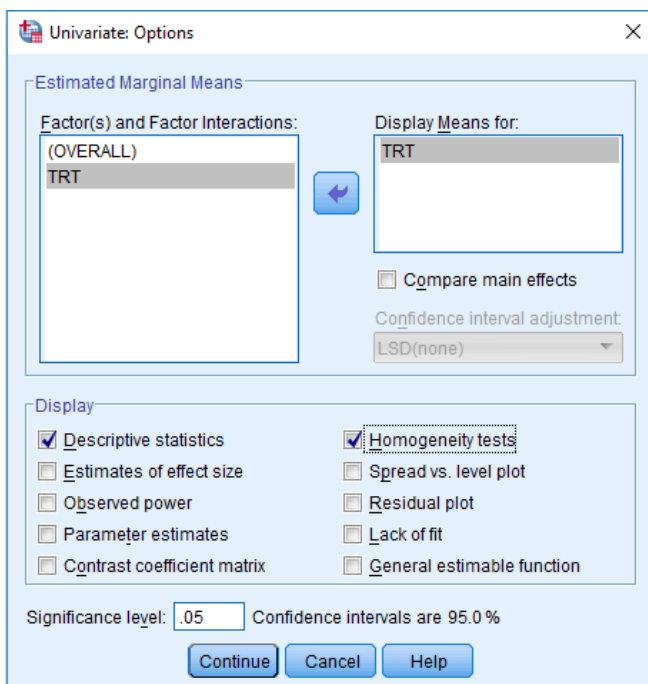
ภาพที่ 7.45 การป้อนข้อมูล TRT ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > General Linear Model > Univariate...** และเลือก ตัวแปร **L*Value [L]** (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร **Treatment [TRT]** (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 7.46



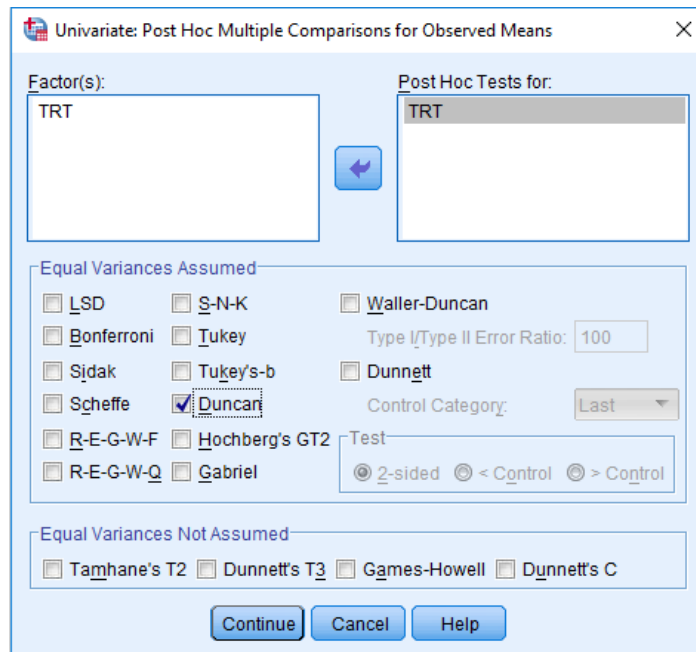
ภาพที่ 7.46 การกำหนดตัวแปรในเมนูคำสั่ง Univariate

(3) เลือกคำสั่ง **Options...** จะได้หน้าจอตั้งภาพที่ 7.47 กดเลือก ตัวแปร **TRT** ใส่ในกล่อง **Display Means for** ที่อยู่ด้านขวามือ และกดเลือก **Homogeneity tests** (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก **Descriptive statistics** (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด **Continue** เพื่อกลับมายังหน้าจอหลักของคำสั่ง **Univariate**



ภาพที่ 7.47 การกำหนดตัวแปรในเมนูคำสั่ง Univariate : Options

(4) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 7.48 ให้กดเลือก **ตัวแปร TRT** ในกล่อง Post Hoc Tests for: และกดเลือก **วิธี Duncan** ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 7.49



ภาพที่ 7.48 การกำหนดตัวแปรและวิธีการทดสอบในเมนูคำสั่ง Post Hoc...

Univariate Analysis of Variance

Between-Subjects Factors			
		Value Label	N
Treatment	1	Taro25%Sugar0%	3
	2	Taro50%Sugar0%	3
	3	Taro25%Sugar25%	3
	4	Taro50%Sugar25%	3

Descriptive Statistics			
Dependent Variable: L*Value			
Treatment	Mean	Std. Deviation	N
Taro25%Sugar0%	57.6933	.17616	3
Taro50%Sugar0%	41.5833	.32517	3
Taro25%Sugar25%	61.5533	.21032	3
Taro50%Sugar25%	66.4400	.39737	3
Total	56.8175	9.74344	12

Levene's Test of Equality of Error Variances ^a			
Dependent Variable: L*Value			
F	df1	df2	Sig.
.764	3	8	.545

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + TRT

Tests of Between-Subjects Effects					
Dependent Variable: L*Value					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	1043.603 ^a	3	347.868	4105.843	.000
Intercept	38738.740	1	38738.740	457229.149	.000
TRT	1043.603	3	347.868	4105.843	.000
Error	.678	8	.085		
Total	39783.020	12			
Corrected Total	1044.280	11			

a. R Squared = .999 (Adjusted R Squared = .999)

Estimated Marginal Means

Treatment				
Dependent Variable: L*Value				
Treatment	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Taro25%Sugar0%	57.693	.168	57.306	58.081
Taro50%Sugar0%	41.583	.168	41.196	41.971
Taro25%Sugar25%	61.553	.168	61.166	61.941
Taro50%Sugar25%	66.440	.168	66.052	66.828

Post Hoc Tests

Treatment

Homogeneous Subsets

L*Value					
Duncan ^{a, b}					
Treatment	N	Subset			
		1	2	3	4
Taro50%Sugar0%	3	41.5833			
Taro25%Sugar0%	3		57.6933		
Taro25%Sugar25%	3			61.5533	
Taro50%Sugar25%	3				66.4400
Sig.		1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .085.

a. Uses Harmonic Mean Sample Size = 3.000.

b. Alpha = .05.

ภาพที่ 7.49 ผลลัพธ์การวิเคราะห์ความแปรปรวนตัวแปร TRT และ ตัวแปร L

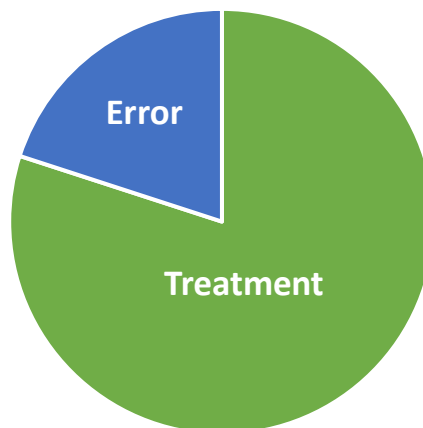


บทที่ 8

การวางแผนการทดลองแบบสุ่มสมบูรณ์

การวางแผนการทดลองแบบสุ่มสมบูรณ์ (Completely Randomized Design : CRD) เป็นการจัดสิ่งทดลองหรือ ทรีทเมนต์ (Treatment) ให้กับหน่วยทดลอง (Experimental unit) โดยสุ่มอย่างสมบูรณ์และไม่มีการกำหนดเงื่อนไขใดๆ เกี่ยวกับการสุ่มสิ่งทดลอง ดังนั้นหน่วยทดลองทุกหน่วยจะมีโอกาสได้รับสิ่งทดลองใดสิ่งทดลองหนึ่งโดยเท่าเทียมกัน การวางแผนการทดลองแบบ CRD นี้จึงเหมาะสำหรับงานวิจัยที่นักวิจัยมั่นใจว่าหน่วยทดลองทุกหน่วยมีความสม่ำเสมอทั้งหมด ไม่มีความแตกต่างเนื่องจากปัจจัยอื่นๆ เช่น อายุ น้ำหนัก

การวางแผนการทดลองแบบ CRD มีการผันแปรของสิ่งทดลองเพียงทางเดียวหรือปัจจัยเดียว โดยสิ่งทดลองนั้นมักมีมากกว่า 2 สิ่งทดลองขึ้นไป แหล่งของความแปรปรวนในการวางแผนการทดลองแบบ CRD มี 2 แหล่ง คือ จากสิ่งทดลอง (Treatment) และ จากความคลาดเคลื่อนจากการทดลอง (Experimental error) ดังภาพที่ 8.1 ดังนั้นในการวิเคราะห์ความแปรปรวนสำหรับการทดลองแบบ CRD จึงเป็น การวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) เนื่องจากสัดส่วนความแปรปรวนที่อธิบายได้ (Explained variance) เป็นความแปรปรวนที่เกิดจากสิ่งทดลองเท่านั้น



ภาพที่ 8.1 การแบ่งสัดส่วนความแปรปรวนของแผนการทดลองแบบสุ่มสมบูรณ์ (CRD)

8.1 ข้อดีและข้อเสียของการวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD)

ข้อดีของการวางแผนการทดลองแบบ CRD

- เป็นการวางแผนการทดลองที่ยืดหยุ่น จำนวนซ้ำของแต่ละสิ่งทดลองไม่จำเป็นต้องเท่ากัน แต่ถ้าไม่มีความจำเป็นควรให้จำนวนซ้ำเท่ากันทุกสิ่งทดลอง
- วิธีการคำนวณง่ายและสะดวก แม้จะมีจำนวนซ้ำของสิ่งทดลองต่างกัน
- เมื่อมีข้อมูลสูญหาย (Missing data) จะกระทบกับความจริงน้อยกว่าการวางแผนแบบอื่น
- การใช้การวางแผนการทดลองแบบ CRD สำหรับงานทดลองขนาดเล็กจะมีความเที่ยงตรงสูงกว่าการวางแผนแบบอื่น

ข้อเสียของการวางแผนการทดลองแบบ CRD

- เนื่องจากไม่มีเงื่อนไขในการจัดสิ่งทดลองให้กับหน่วยทดลอง ทำให้ความคลาดเคลื่อนจากการทดลองไปรวมกับความแปรปรวนของหน่วยทดลองทั้งหมด ยกเว้นความแปรปรวนจากสิ่งทดลอง ดังนั้นหากหน่วยทดลองไม่มีความสม่ำเสมอจะไม่สามารถแยกส่วนแตกต่างนี้ออกจากความคลาดเคลื่อนได้

8.2 ตัวแบบ (Model)

ในการวางแผนการทดลองแบบ CRD นักวิจัยสนใจปัจจัยเพียงหนึ่งปัจจัยที่มีระดับต่างกัน a ระดับ ในบางครั้งระดับของปัจจัย (Levels of factor) เรียกว่า สิ่งทดลองหรือทรีทเมนต์ (Treatment) และในแต่ละสิ่งทดลองมีจำนวนค่าสังเกต (n) เท่ากัน รูปแบบข้อมูลของการทดลองที่มีปัจจัยเดียวแสดงดังตารางที่ 8.1 และตัวแบบ (Model) ของการวางแผนการทดลองแบบ CRD เขียนได้ ดังนี้

$$y_{ij} = \mu + \tau_i + \varepsilon_{ij} ; \quad \begin{array}{l} i = 1, 2, 3, \dots, a \\ j = 1, 2, 3, \dots, n \end{array}$$

โดยที่	y_{ij}	คือ	ค่าสังเกตที่วัดได้จากสิ่งทดลองที่ i หน่วยทดลองที่ j
	μ	คือ	ค่าเฉลี่ยข้อมูลทั้งหมด (Overall mean)
	τ_i	คือ	อิทธิพล (Effect) จากสิ่งทดลองที่ i
	ε_{ij}	คือ	ค่าความคลาดเคลื่อนจากการทดลอง (Experimental error หรือ Random error) ที่อยู่ในค่าสังเกตของสิ่งทดลองที่ i หน่วยทดลองที่ j

ตารางที่ 8.1 แสดงค่าสังเกตของตัวอย่างสุ่มสำหรับการวิเคราะห์ความแปรปรวนของข้อมูลที่ได้จากการวางแผนการทดลองแบบสุ่มสมบูรณ์

สิ่งทดลองที่	ค่าสังเกต						ผลรวม	ค่าเฉลี่ย
	1	2	...	j	...	n		
1	y_{11}	y_{12}	...	y_{1j}	...	y_{1n}	$y_{1\cdot}$	$\bar{y}_{1\cdot}$
2	y_{21}	y_{22}	...	y_{2j}	...	y_{2n}	$y_{2\cdot}$	$\bar{y}_{2\cdot}$
.	.	.	...	.	...	.	.	.
.	.	.	...	.	...	.	.	.
i	y_{i1}	y_{i2}	...	y_{ij}	...	y_{in}	$y_{i\cdot}$	$\bar{y}_{i\cdot}$
.	.	.	...	.	...	.	.	.
.	.	.	...	.	...	.	.	.
a	y_{a1}	y_{a2}	...	y_{aj}	...	y_{an}	$y_{a\cdot}$	$\bar{y}_{a\cdot}$
							$y_{..}$	$\bar{y}_{..}$

โดยที่	y_{ij}	คือ	ค่าสังเกตที่วัดได้จากสิ่งทดลองที่ i หน่วยทดลองที่ j
	$y_{i\cdot}$	คือ	ผลรวมของค่าสังเกตทั้งหมดที่ได้รับสิ่งทดลองที่ i
	$\bar{y}_{i\cdot}$	คือ	ค่าเฉลี่ยของค่าสังเกตทั้งหมดที่ได้รับสิ่งทดลองที่ i
	$y_{..}$	คือ	ผลรวมของค่าสังเกตทั้งหมด
	$\bar{y}_{..}$	คือ	ค่าเฉลี่ยของค่าสังเกตทั้งหมด
	$N = an$	คือ	ผลรวมของจำนวนของค่าสังเกตทั้งหมด

ตัวแบบ (Model) สำหรับการทดสอบสมมติฐานเกี่ยวกับอิทธิพลของสิ่งทดลอง (Treatment effect) แบ่งได้ออกเป็น 2 ประเภท คือ ตัวแบบอิทธิพลคงที่ (Fixed effect model) และตัวแบบอิทธิพลสุ่ม (Random effect model) ในการวิเคราะห์ความแปรปรวนนี้นักวิจัยต้องการศึกษาปัจจัยที่คาดว่าจะทำให้ตัวแปรเชิงปริมาณมีค่าเฉลี่ยต่างกัน ความแตกต่างของตัวแบบอิทธิพลคงที่และตัวแบบอิทธิพลสุ่ม มีดังนี้

(1) **ตัวแบบอิทธิพลคงที่ (Fixed effect model)** คือ ระดับของปัจจัยหรือสิ่งทดลองเป็นค่าคงที่ที่นักวิจัยกำหนดขึ้นเอง ดังนั้นเมื่อทดสอบสมมติฐานเกี่ยวกับค่าเฉลี่ยของสิ่งทดลอง ผลสรุปจากการวิเคราะห์จะสามารถใช้สรุปได้เพียงระดับต่างๆ ตามที่นักวิจัยกำหนดไว้เท่านั้น ไม่สามารถขยายผลไปยังสิ่งทดลองที่คล้ายคลึงกันได้ โดยทั่วไปแล้วในงานวิจัยและพัฒนาผลิตภัณฑ์นักวิจัยมักกำหนดระดับของปัจจัยที่สนใจขึ้นเอง จึงใช้ตัวแบบอิทธิพลคงที่ในการประมาณค่าพารามิเตอร์สำหรับทดสอบสมมติฐานงานวิจัย ตัวอย่างการทดลอง

ที่มีตัวแบบอิทธิพลคงที่ เช่น นักวิจัยศึกษาอิทธิพลของน้ำมันพืช 3 ชนิดที่มีผลต่อการอมน้ำมันของข้าวเกรียบปลาทุ ได้แก่ น้ำมันถั่วเหลือง, น้ำมันปาล์ม และ น้ำมันรำข้าว ในการศึกษาครั้งนี้จะมี 3 สิ่งทดลอง (น้ำมันพืช 3 ชนิด) ซึ่งนักวิจัยระบุแบบเฉพาะเจาะจง

(2) **ตัวแบบอิทธิพลสุ่ม (Random effect model)** คือ ระดับของปัจจัยหรือสิ่งทดลองเป็นตัวอย่างของระดับที่เป็นไปได้ทั้งหมด ดังนั้นกรณีนี้นักวิจัยสามารถสรุปผลโดยการขยายไปยังสิ่งทดลองทั้งหมดในประชากรได้ไม่ว่าระดับของสิ่งทดลองนั้นจะถูกนำมาใช้ในการทดลองหรือไม่ก็ตาม ตัวอย่างการทดลองที่มีตัวแบบอิทธิพลสุ่ม เช่น นักวิจัยศึกษาอิทธิพลของน้ำมันถั่วเหลืองที่มีผลต่อการอมน้ำมันของข้าวเกรียบปลาทุ จึงได้สุ่มน้ำมันถั่วเหลืองทางการค้ามา 3 ยี่ห้อ นำมาใช้ทอดข้าวเกรียบปลาทุแล้วนำไปวิเคราะห์ปริมาณไขมันในห้องปฏิบัติการ ในกรณีนี้สิ่งทดลองมี 3 สิ่งทดลอง (ยี่ห้อน้ำมันถั่วเหลือง) ซึ่งทั้ง 3 ยี่ห้อถูกสุ่มมาแบบไม่เฉพาะเจาะจงจากน้ำมันถั่วเหลืองทางการค้าทั้งหมดที่จำหน่ายในท้องตลาด

ในหนังสือเล่มนี้จะอธิบายเฉพาะการใช้ตัวแบบอิทธิพลคงที่ในการประมาณค่าพารามิเตอร์สำหรับทดสอบสมมติฐานงานวิจัย

8.3 การตั้งสมมติฐานเพื่อการทดสอบความสัมพันธ์หรือทดสอบเปรียบเทียบค่าเฉลี่ย

สำหรับการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบสุ่มสมบูรณ์ประกอบด้วย การทดสอบสมมติฐานด้วยวิธีการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) หรือเรียกว่า การวิเคราะห์ความแปรปรวนแบบสุ่มสมบูรณ์ (One-Way Completely Randomized Design) ซึ่งการเขียนสมมติฐานทางสถิติ สามารถเขียนได้ ดังนี้

H_0 : ค่าเฉลี่ยของตัวอย่างทุกสิ่งทดลองไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\mu_1 = \mu_2 = \mu_3$ (กรณีมี 3 สิ่งทดลอง)

H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่; $i \neq j$

ถ้ายอมรับสมมติฐานหลัก (Accept H_0) แสดงว่า สิ่งทดลองมีค่าเฉลี่ย μ ไม่แตกต่างกัน หรือ สิ่งทดลองไม่มีอิทธิพลต่อค่าสังเกต

ถ้าปฏิเสธสมมติฐานหลัก (Reject H_0) แสดงว่า สิ่งทดลองมีค่าเฉลี่ย μ แตกต่างกัน หรือ สิ่งทดลองมีอิทธิพลต่อค่าสังเกต ต้องทำการเปรียบเทียบความแตกต่างของค่าเฉลี่ยโดยใช้วิธีการทดสอบพหุคูณ (Multiple

comparison tests) หรือ Post hoc tests กรณีที่ตัวแปรต้นหรือระดับของปัจจัยที่สนใจมี k กลุ่ม จะต้องทดสอบทั้งหมด kC_2 คู่ เช่น มี 3 อาชีพ ($k=3$) จะมี 3C_2 นั่นเอง

เงื่อนไขของการวิเคราะห์ความแปรปรวน คือ การสุ่มตัวอย่างจะต้องสุ่มอย่างเป็นอิสระกัน, สุ่มตัวอย่างจากประชากรที่มีการแจกแจงแบบปกติ และ ความแปรปรวนของแต่ละประชากรต้องเท่ากัน ถ้าข้อมูลที่น่าสนใจมาวิเคราะห์ทางสถิติเป็นไปตามเงื่อนไขดังกล่าว จะใช้ค่าสถิติทดสอบ F ในตาราง ANOVA สำหรับการทดสอบความสัมพันธ์หรือทดสอบเปรียบเทียบค่าเฉลี่ย ผลการวิเคราะห์ความแปรปรวนแบบทางเดียวอาจสรุปเป็นตารางที่เรียกว่า ตารางวิเคราะห์ความแปรปรวนทางเดียว (One-Way Analysis of Variance Table หรือ ANOVA Table) ดังตารางที่ 8.2

ตารางที่ 8.2 การวิเคราะห์ความแปรปรวนแบบทางเดียวสำหรับการวางแผนการทดลองแบบ CRD

แหล่งความแปรปรวน (Source of Variation)	องศาความเป็นอิสระ(Degree of Freedom)	ผลบวกกำลังสอง (Sum of Squares)	ค่ากำลังสองเฉลี่ย (Mean Squares)	F
ระหว่างสิ่งทดลอง (Between Treatment)	$a - 1$	$\sum_{i=1}^a \frac{y_i^2}{n_i} - \frac{y_{..}^2}{N}$	$\frac{SSTr}{a - 1}$	$\frac{MSTr}{MSE}$
ภายในสิ่งทดลองหรือความคลาดเคลื่อน (Within Treatment or Error)	$N - a$	$SSE = SST - SSTr$	$\frac{SSE}{N - a}$	
ยอดรวม (Total)	$N - 1$	$\sum_{i=1}^a \sum_{j=1}^n y_{ij}^2 - \frac{y_{..}^2}{N}$		

a = จำนวนสิ่งทดลองหรือระดับของปัจจัย, N = ผลรวมของจำนวนของค่าสังเกตทั้งหมด

8.4 ตัวอย่างรูปแบบข้อมูลจากการวางแผนการทดลองแบบสุ่มสมบูรณ์

นักวิจัยต้องการศึกษาอิทธิพลของอุณหภูมิในการทอด 3 ระดับ (140, 150 และ 160 °C) ที่มีผลต่อค่าความแข็งของโดนัท กรณีนี้มีสิ่งทดลอง 3 สิ่งทดลองตามระดับของปัจจัย นั่นคือ อุณหภูมิในการทอด นักวิจัยวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) ในแต่ละสิ่งทดลองจะทำการวัดค่าความแข็ง (Hardness) ของโดนัทจำนวน 5 ชิ้น ด้วยเครื่อง Texture analyzer โดยใช้วิธี Single hardness test

ดังนั้นการทดลองนี้ต้องอยู่ภายใต้เงื่อนไขที่ว่า แป้งโดนัทที่เอามาทอดต้องมีความสม่ำเสมอ เช่น เป็นแป้งโดนัทที่ได้มาจากการผสมส่วนผสมสูตรเดียวกันในกระบวนการผสมเดียวกัน ดังนั้นเมื่อส่วนผสมผสมของโดนัทเสร็จแล้วจะทำการแบ่งและขึ้นรูปเป็นโดนัท 15 ชิ้นเพื่อเอามาทอดที่อุณหภูมิต่างกัน โดยสุ่มมาทอดอุณหภูมิละ 5 ชิ้น ดังนั้น จะมีหน่วยทดลอง (โดนัท) ทั้งหมด 15 หน่วย ซึ่งสามารถเขียนรูปแบบข้อมูลได้ ดังตารางที่ 8.3



ตารางที่ 8.3 ผังแสดงหน่วยทดลอง (โดนัท) ที่ถูกสุ่มในแต่ละสิ่งทดลองเมื่อวางแผนการทดลองแบบ CRD

สิ่งทดลองที่ (อุณหภูมิทอด)	ค่าความแข็ง (ซ้ำที่/ ชิ้นที่)				
	1	2	3	4	5
สิ่งทดลองที่ 1 อุณหภูมิ 140 °C					
สิ่งทดลองที่ 2 อุณหภูมิ 150 °C					
สิ่งทดลองที่ 3 อุณหภูมิ 160 °C					

จากตารางที่ 8.3 ตัวเลขที่อยู่บนโดนัท คือ เลขของแป้งโดนัทที่ยังไม่ได้ทอดทั้ง 15 ชิ้นและถูกสุ่มมาทอดอุณหภูมิละ 5 ชิ้น เช่น แป้งโดนัทชิ้นที่ 11, 2, 6, 14 และ 1 จะถูกนำมาทอดที่อุณหภูมิ 140 °C

จากตัวอย่างข้างต้นจึงสามารถสรุปรายละเอียดของข้อมูลได้ ดังนี้

- (1) ปัจจัยที่ศึกษา มี 1 ปัจจัย คือ อุณหภูมิในการทอด
- (2) จำนวนสิ่งทดลอง (Treatment) เท่ากับ 3 สิ่งทดลอง (ตามระดับของอุณหภูมิในการทอด)
- (3) แต่ละสิ่งทดลอง มีหน่วยทดลอง 5 หน่วย หรือมีจำนวนซ้ำ (Replication) เท่ากับ 5 ซ้ำ
- (4) จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 15 หน่วย (โดนัท 15 ชิ้น)
- (5) ค่าสังเกต (ตัวแปรตาม) คือ ค่าความแข็ง (Hardness)

8.5 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD

การวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) คือ การวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) ดังนั้นคำสั่งที่ใช้วิเคราะห์ข้อมูล คือ

Analyze ➤ General Linear Model ➤ Univariate...

8.5.1 ตัวอย่างข้อมูล

นักวิจัยต้องการทราบว่าอุณหภูมิที่ใช้ในการอบมีผลต่อค่าสีของผักแผ่นหรือไม่ จึงวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) โดยการนำส่วนผสมของผักแผ่นที่เตรียมไว้มาแบ่งออกเป็น 4 ส่วน แล้วนำไปอบด้วยอุณหภูมิที่ต่างกัน 4 ระดับ คือ 50, 60, 70 และ 80 °C โดยใช้ตู้อบลมร้อนเป็นเวลา 3 ชั่วโมงเท่ากัน นำตัวอย่างผักแผ่นหลังจากการอบไปวิเคราะห์ค่าสี L* ด้วยเครื่อง Spectrophotometer โดยนำตัวอย่างมาบดก่อนเพื่อลดผลกระทบจากลักษณะของผักแผ่นที่ไม่สม่ำเสมอ จากนั้นทำการบรรจุตัวอย่างผักแผ่นผงในภาชนะบรรจุตัวอย่างแล้ววิเคราะห์ค่าสี L* สิ่งทดลองละ 6 ซ้ำ ได้ผลการวิเคราะห์แสดงดังตารางที่ 8.4

ตารางที่ 8.4 ค่าสี L* ของผักแผ่นที่อบด้วยอุณหภูมิที่ต่างกัน

การวัดค่าซ้ำที่	ค่าสี L* ที่ใช้อุณหภูมิในการอบ (°C)			
	50	60	70	80
1	67.4	56.6	54.4	50.2
2	66.5	57.2	53.5	51.6
3	65.5	56.1	54.5	49.3
4	68.9	57.2	55.6	50.9
5	65.4	58.8	56.3	49.5
6	66.3	55.3	57.7	51.1

8.5.2 การเขียนสมมติฐานเพื่อทดสอบ

จากข้อมูลดังกล่าว สามารถสรุปรายละเอียดของข้อมูลได้ ดังนี้

- (1) ปัจจัยที่ศึกษา มี 1 ปัจจัย คือ อุณหภูมิในการอบ
- (2) จำนวนสิ่งทดลอง (Treatment) เท่ากับ 4 สิ่งทดลอง (ตามระดับของอุณหภูมิในการอบ)
- (3) แต่ละสิ่งทดลอง มีหน่วยทดลอง 6 หน่วย หรือมีจำนวนซ้ำ (Replication) เท่ากับ 6 ซ้ำ
- (4) จำนวนหน่วยทดลอง (Experimental unit) เท่ากับ 24 หน่วย
- (5) คำสั่งเกต (ตัวแปรตาม) คือ ค่าสี L*

นักวิจัยสามารถเขียนสมมติฐานทางสถิติสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ ดังนี้

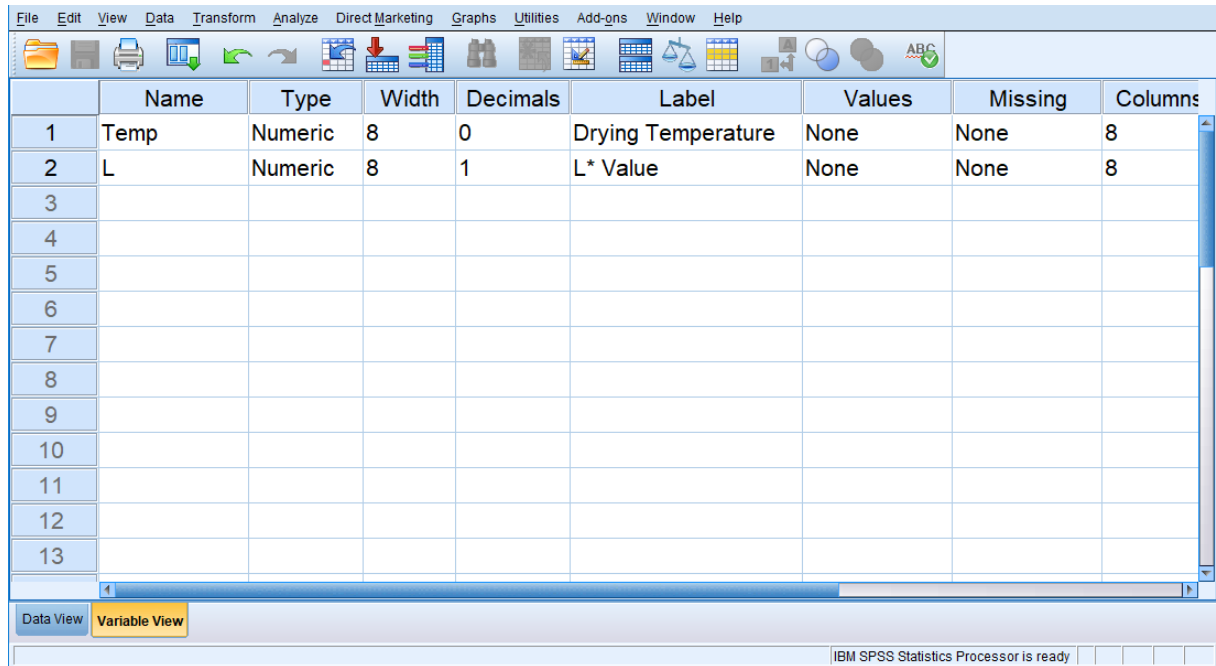
- หรือ
- H_0 : ค่าเฉลี่ยสี L^* ของผักแผ่นไม่ขึ้นอยู่กับอุณหภูมิที่ใช้ในการอบ
- H_1 : ค่าเฉลี่ยสี L^* ของผักแผ่นขึ้นกับอุณหภูมิที่ใช้ในการอบ
- หรือ
- H_0 : อุณหภูมิที่ใช้ในการอบไม่มีอิทธิพลต่อค่าเฉลี่ยสี L^* ของผักแผ่น
- H_1 : อุณหภูมิที่ใช้ในการอบมีอิทธิพลต่อค่าเฉลี่ยสี L^* ของผักแผ่น
- หรือ
- H_0 : ค่าเฉลี่ยสี L^* ของผักแผ่นทุกสิ่งทดลองไม่แตกต่างกัน
- H_1 : ค่าเฉลี่ยสี L^* ของผักแผ่นอย่างน้อย 2 สิ่งทดลองแตกต่างกัน
- หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้
- H_0 : $\mu_1 = \mu_2 = \mu_3 = \mu_4$
- H_1 : $\mu_i \neq \mu_j$ อย่างน้อย 1 คู่; $i \neq j$

8.5.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลค่าสี L^* ของผักแผ่นที่อบด้วยอุณหภูมิที่แตกต่างกัน จะต้องใช้ตัวแปร var จำนวน 2 ตัวแปรสำหรับป้อนข้อมูล 2 ตัว ได้แก่ อุณหภูมิในการอบ และ ค่าสี L^* ดังนั้น กำหนดชื่อตัวแปรเป็น **Temp** และ **L** ในหน้าต่าง Variable View ดังภาพที่ 8.2 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
อุณหภูมิในการอบ	Temp	Drying Temperature	0	ไม่ต้องกำหนด
ค่าสี L^*	L	L^* Value	1	ไม่ต้องกำหนด

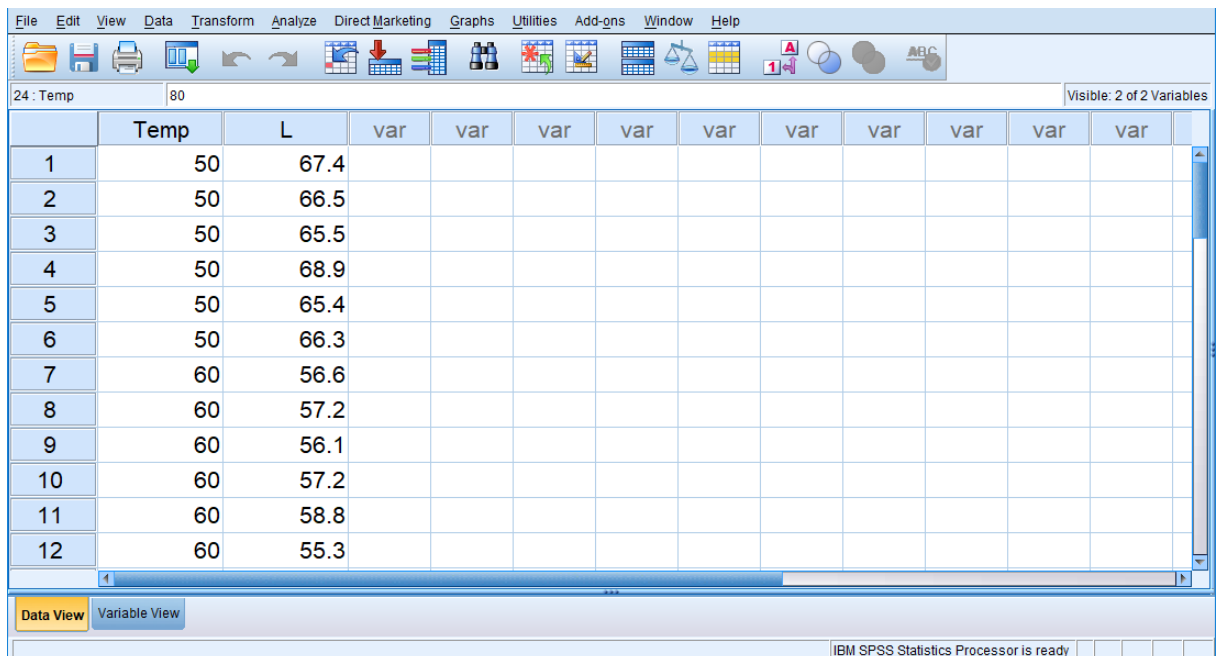
หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



	Name	Type	Width	Decimals	Label	Values	Missing	Columns
1	Temp	Numeric	8	0	Drying Temperature	None	None	8
2	L	Numeric	8	1	L* Value	None	None	8
3								
4								
5								
6								
7								
8								
9								
10								
11								
12								
13								

ภาพที่ 8.2 กำหนดรายละเอียดตัวแปร Temp และ L ในหน้าต่าง Variable View

(2) ป้อนข้อมูลค่าสี L* เรียงตามอุณหภูมิในการอบในหน้าต่าง Data View ดังภาพที่ 8.3

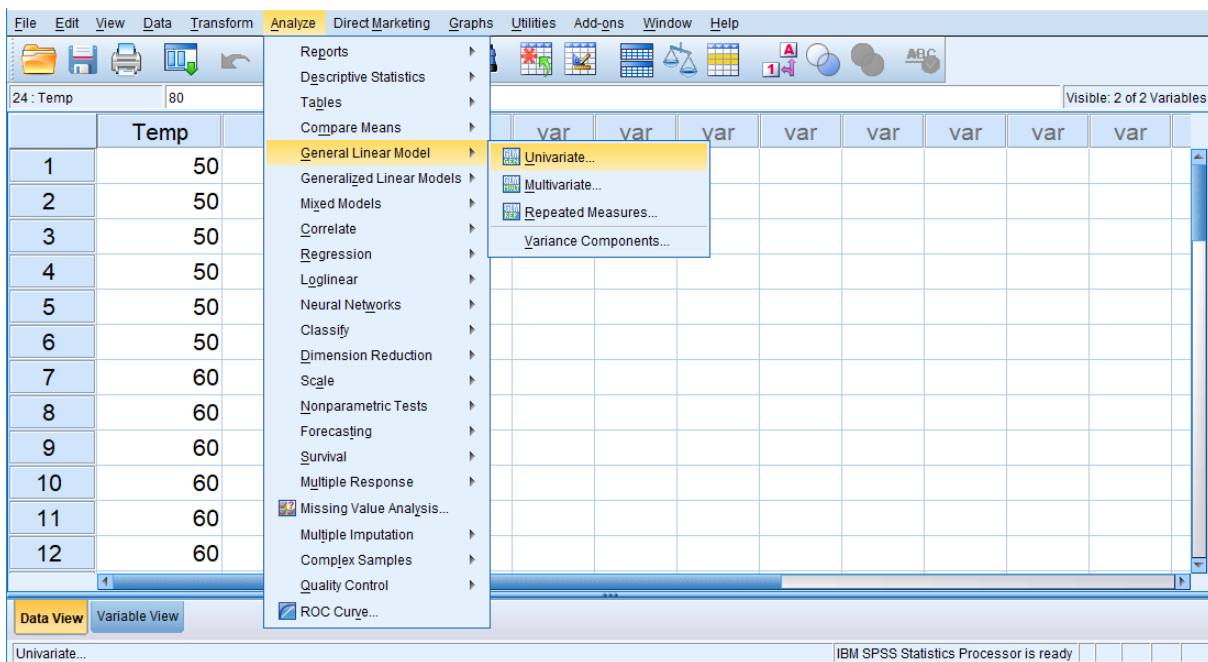


	Temp	L	var	var	var	var	var	var	var	var	var	var
1	50	67.4										
2	50	66.5										
3	50	65.5										
4	50	68.9										
5	50	65.4										
6	50	66.3										
7	60	56.6										
8	60	57.2										
9	60	56.1										
10	60	57.2										
11	60	58.8										
12	60	55.3										

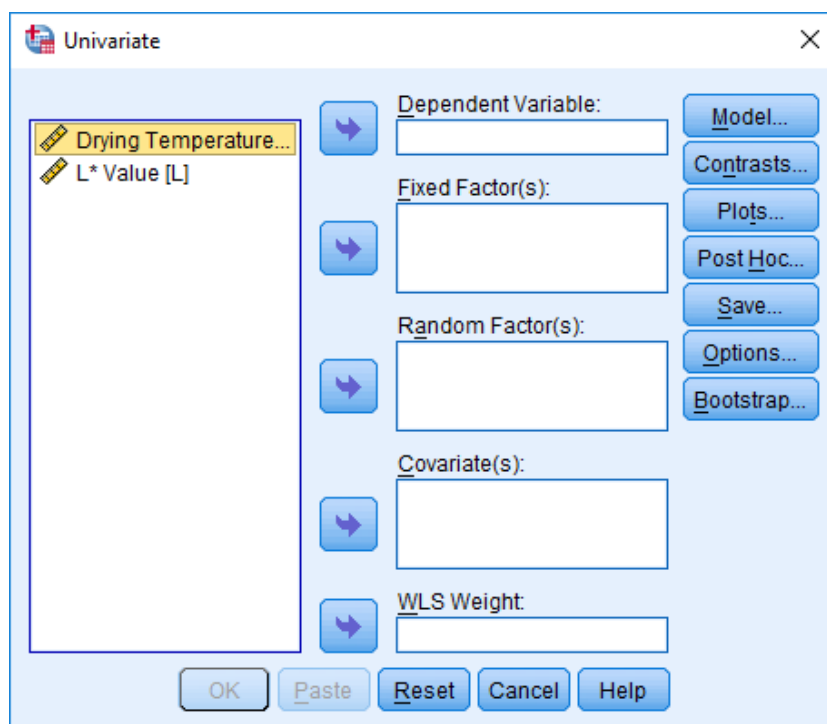
ภาพที่ 8.3 ป้อนข้อมูลตัวแปร Temp และ L ลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 8.4 จะ

ปรากฏหน้าต่างดังภาพที่ 8.5

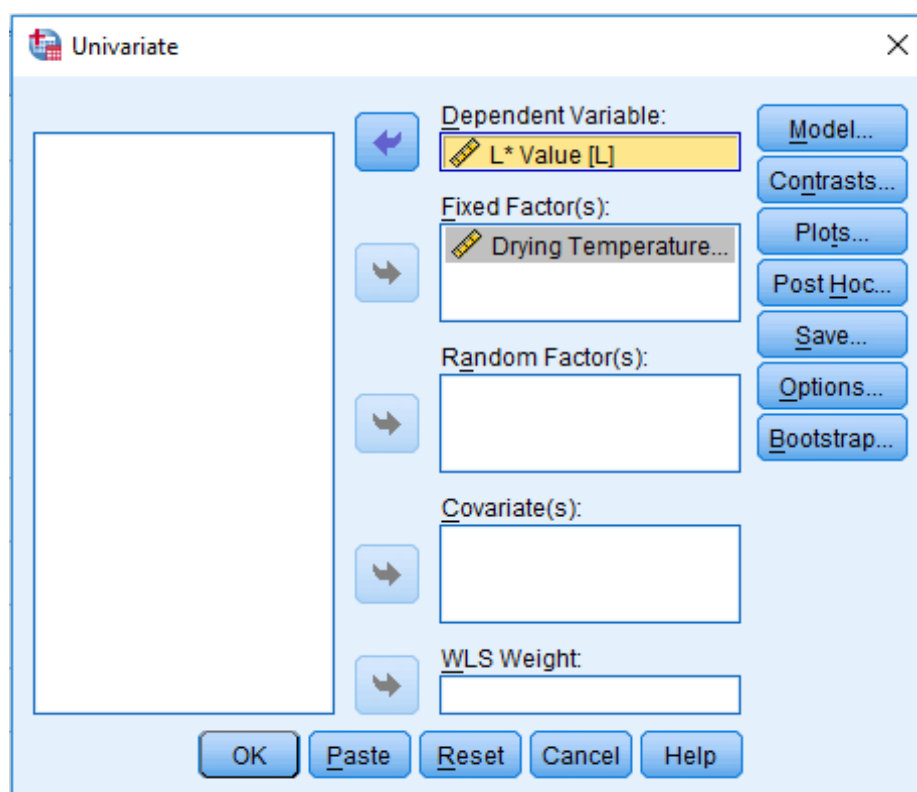


ภาพที่ 8.4 เมนูหน้าต่างคำสั่ง Analyze > General Linear Model > Univariate...



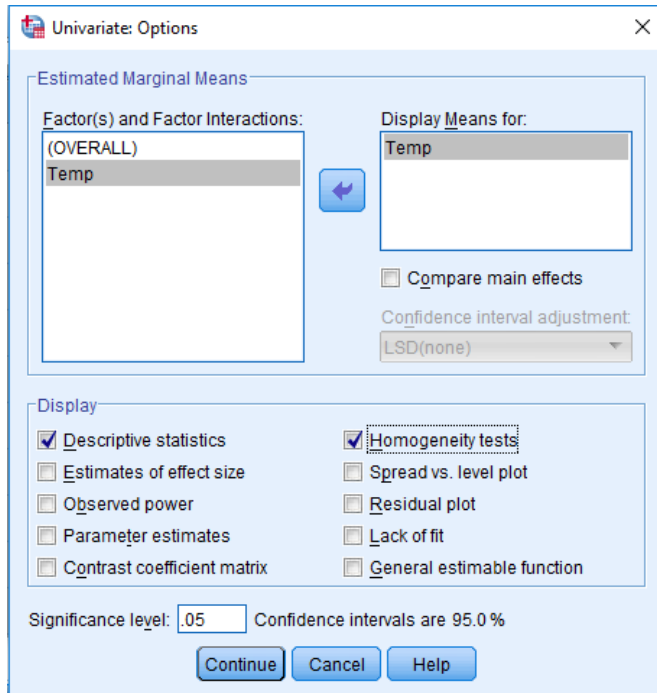
ภาพที่ 8.5 หน้าจอคำสั่ง Univariate

(4) เลือก ตัวแปร **L* Value [L]** (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร **Drying Temperature [Temp]** (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 8.6



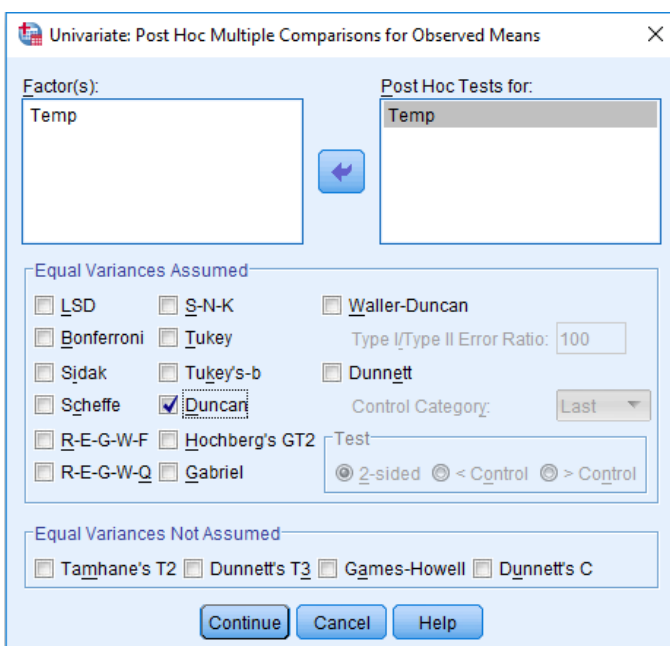
ภาพที่ 8.6 เลือกตัวแปร L* Value [L] ใส่ในกล่อง Dependent Variable และ ตัวแปร Drying Temperature [Temp] ใส่ในกล่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Options...** จะได้หน้าจอ ดังภาพที่ 8.7 ให้เลือก **ตัวแปร Temp** จากกล่องซ้ายมือ ไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลในแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 8.7 หน้าจอคำสั่ง Univariate : Options

(6) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือ เปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอ ดังภาพที่ 8.8 ให้กดเลือก **ตัวแปร Temp** จาก กล่องซ้ายมือไปใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธี อื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 8.9



ภาพที่ 8.8 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

Univariate Analysis of Variance

Between-Subjects Factors		
		N
Drying Temperature	50	6
	60	6
	70	6
	80	6

Descriptive Statistics			
Dependent Variable: L* Value			
Drying Temperature	Mean	Std. Deviation	N
50	66.667	1.3155	6
60	56.867	1.1894	6
70	55.333	1.5188	6
80	50.433	.9201	6
Total	57.325	6.1330	24

Levene's Test of Equality of Error Variances ^a			
Dependent Variable: L* Value			
F	df1	df2	Sig.
.467	3	20	.709

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Temp

Tests of Between-Subjects Effects					
Dependent Variable: L* Value					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	833.632 ^a	3	277.877	176.467	.000
Intercept	78867.735	1	78867.735	50085.352	.000
Temp	833.632	3	277.877	176.467	.000
Error	31.493	20	1.575		
Total	79732.860	24			
Corrected Total	865.125	23			

a. R Squared = .964 (Adjusted R Squared = .958)

Estimated Marginal Means

Drying Temperature				
Dependent Variable: L* Value				
Drying Temperature	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
50	66.667	.512	65.598	67.735
60	56.867	.512	55.798	57.935
70	55.333	.512	54.265	56.402
80	50.433	.512	49.365	51.502

Post Hoc Tests

Drying Temperature

Homogeneous Subsets

L* Value					
Duncan ^{a, b}					
Drying Temperature	N	Subset			
		1	2	3	4
80	6	50.433			
70	6		55.333		
60	6			56.867	
50	6				66.667
Sig.		1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = 1.575.

a. Uses Harmonic Mean Sample Size = 6.000.

b. Alpha = .05.

ภาพที่ 8.9 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนด้วยคำสั่ง Univariate

8.5.4 การอ่านผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ CRD

จากผลการวิเคราะห์ในภาพที่ 8.9 แสดงผลลัพธ์ได้ 6 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Between-Subjects Factors เป็นส่วนที่แสดงรายละเอียดของข้อมูลที่ป้อน ดังนี้

Drying Temperature คือ ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 4 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการอบ ได้แก่ 50, 60, 70 และ 80 °C

N คือ จำนวนข้อมูลในแต่ละสิ่งทดลอง ในที่นี้แต่ละสิ่งทดลองมีจำนวนข้อมูลเท่ากับ 6

ส่วนที่ 2 ตาราง Descriptive Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบจากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Drying Temperature	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 4 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการอบ ได้แก่ 50, 60, 70 และ 80 °C
Mean	คือ	ค่าเฉลี่ย L^* ของแต่ละสิ่งทดลองจำแนกโดยระดับอุณหภูมิในการอบ
Std. Deviation	คือ	ค่าส่วนเบี่ยงเบนมาตรฐานของค่า L^* ที่แสดงการกระจายของข้อมูล
N	คือ	จำนวนข้อมูลในแต่ละสิ่งทดลอง ในที่นี้แต่ละสิ่งทดลองมีจำนวนข้อมูลเท่ากัน คือ 6

ส่วนที่ 3 ตาราง Levene's Test of Equality of Error Variances ใช้พิจารณาการกระจายของข้อมูลหรือความแปรปรวนของข้อมูลแต่ละกลุ่ม อธิบายได้ ดังนี้

F, df1, df2	คือ	ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในกรณีนี้การทดสอบสมมติฐานเป็นแบบสองทาง (Two-tailed test) เขียนได้ดังนี้ H_0 : ข้อมูลทั้ง 4 กลุ่มมีการกระจายตัวไม่แตกต่างกัน H_1 : ข้อมูลทั้ง 4 กลุ่มมีการกระจายตัวแตกต่างกัน ค่า Sig. มีค่าเท่ากับ 0.709 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 4 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ ระดับนัยสำคัญ 0.05 (Equal variance assumed)

ส่วนที่ 4 ตาราง Tests of Between-Subjects Effects เป็นส่วนที่แสดงค่าสถิติต่างๆ ของตารางวิเคราะห์ความแปรปรวน เพื่อใช้ในการทดสอบสมมติฐานทางสถิติด้วยค่าสถิติ F หรือค่าความน่าจะเป็น Sig. เพื่อใช้ทดสอบสมมติฐาน อธิบายได้ ดังนี้

Type III Sum of Squares, df, Mean Square, F	คือ	ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในที่นี้จะพิจารณาเฉพาะค่า Sig. ของตัวแปร Temp เท่านั้น ตามสมมติฐานที่เขียนไว้ ดังนี้ H_0 : ค่าเฉลี่ย L^* ของผักแผ่นทุกสิ่งทดลองไม่แตกต่างกัน H_1 : ค่าเฉลี่ย L^* ของผักแผ่นอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

ค่า Sig. มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ย L^* ของผักแผ่นอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ต้องพิจารณาความแตกต่างของค่าเฉลี่ยแต่ละสิ่งทดลองต่อไปในส่วนที่ 6

ส่วนที่ 5 ตาราง Estimated Marginal Means เป็นส่วนที่แสดงค่าเฉลี่ยโดยประมาณของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Drying Temperature	คือ ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 4 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการอบ ได้แก่ 50, 60, 70 และ 80 °C
Mean	คือ ค่าเฉลี่ย L^* ของแต่ละสิ่งทดลองโดยประมาณ
Std. Error	คือ ค่าความคลาดเคลื่อนมาตรฐานในแต่ละกลุ่ม
95% Confidence Interval	คือ ค่าที่แสดงขอบเขตบนล่างของค่าเฉลี่ยในแต่ละกลุ่มในช่วงความเชื่อมั่น 95%

ส่วนที่ 6 ตาราง Post Hoc Tests เป็นส่วนที่แสดงค่าสถิติสำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Test) ซึ่งส่วนนี้จะนำมาใช้พิจารณาต่อเมื่อในส่วนที่ 4 ตาราง Tests of Between-Subjects Effects มีการปฏิเสธสมมติฐานหลัก (H_0) อธิบายได้ ดังนี้

Drying Temperature	คือ ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 4 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการอบ ได้แก่ 50, 60, 70 และ 80 °C
N	คือ จำนวนข้อมูลในแต่ละสิ่งทดลอง ในที่นี้แต่ละสิ่งทดลองมีจำนวนข้อมูลเท่ากัน คือ 6
Subset	คือ แสดงการแบ่งกลุ่มของสิ่งทดลอง ให้พิจารณารายละเอียดของ Subset ดังนี้ (1) มีจำนวน 4 Subset แต่ละ Subset มีสมาชิก 1 สิ่งทดลอง นั้นแสดงว่าทุกสิ่งทดลองมีค่าเฉลี่ย L^* แตกต่างกันอย่างมีนัยสำคัญที่ระดับ 0.05 (2) ตัวเลขที่แสดงในแต่ละ Subset คือ ค่าเฉลี่ย L^* ของระดับอุณหภูมินั้นๆ เช่น Subset ที่ 1 ตัวเลข 50.433 แสดงค่าเฉลี่ย L^* ของอุณหภูมิ 80 °C (3) แต่ละ Subset จะใช้ตัวอักษรภาษาอังกฤษพิมพ์เล็กในการบ่งชี้กลุ่ม โดยจะนำตัวอักษรนี้ไปแสดงหลัง ค่าเฉลี่ย $\pm$ ค่าส่วนเบี่ยงเบนมาตรฐาน เช่น $50.43 \pm 0.92d$ กรณีนี้มี 4 Subset จึงใช้ตัวอักษร a, b, c และ d บ่งชี้สมาชิกที่อยู่ใน Subset ที่ 4, 3, 2 และ 1 ตามลำดับ (เรียงลำดับตัวอักษรตามค่าเฉลี่ยจากมากไปน้อย ทั้งนี้ นักวิจัยอาจเรียงจากน้อยไปมากได้เช่นกัน)

8.5.5 การนำเสนอผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ CRD

จากภาพที่ 8.9 ในตาราง Tests of Between-Subjects Effects มีการปฏิเสธสมมติฐานหลัก (H_0) จึงสรุปได้ว่า ค่าเฉลี่ยสี L^* ของผักแผ่นอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 จึงทำการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Test) ด้วยวิธี Duncan แล้วพบว่า ค่าเฉลี่ยสี L^* ของผักแผ่นทุกระดับอุณหภูมิอบแตกต่างกันที่ระดับนัยสำคัญ 0.05 (เนื่องจากอยู่กันคนละ Subset) โดยในการเขียนผลการวิเคราะห์ในรูปแบบตารางจะนำข้อมูลจากส่วนที่ 2 ตาราง Descriptive Statistics และส่วนที่ 6 ตาราง Post Hoc Tests มาใช้ในการเขียนผล ดังแสดงในตารางที่ 8.5

ตารางที่ 8.5 ค่าเฉลี่ยสี L^* ของผักแผ่นที่ใช้อุณหภูมิในการอบแตกต่างกัน

อุณหภูมิในการอบ ($^{\circ}\text{C}$)	ค่าเฉลี่ยสี L^*
50	$66.67 \pm 1.32a$
60	$56.87 \pm 1.19b$
70	$55.33 \pm 1.52c$
80	$50.43 \pm 0.92d$

^{a-d} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

8.6 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD

8.6.1 ตัวอย่างข้อมูล

นิสิตปริญญาตรีคนหนึ่งต้องการศึกษาผลของการใช้แป้งฟักทองแทนที่แป้งสาลีบางส่วนต่อลักษณะเนื้อสัมผัสของชิฟฟอนเค้ก จึงได้ทำการแทนที่แป้งสาลีด้วยแป้งฟักทอง 3 ระดับ คือ 0, 10 และ 20% โดยวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) นำชิฟฟอนเค้กที่ได้แต่ละสิ่งทดลองมาหั่นให้มีขนาด $20 \times 20 \times 20$ ลูกบาศก์มิลลิเมตร แล้วนำไปวิเคราะห์ค่าเนื้อสัมผัสด้วยวิธี Texture profile analysis โดยใช้เครื่อง Texture analyzer ทำการวัดค่า 5 ข้อ บันทึกค่าความแข็ง (Hardness, N), ค่าความสามารถในการเกาะติดกัน (Cohesiveness), ค่าความยืดหยุ่น (Springiness, mm) และค่าความยากง่ายในการเคี้ยว (Chewiness, Nm) ได้ผลของการวิเคราะห์ค่าเนื้อสัมผัสแสดงดังตารางที่ 8.6

ตารางที่ 8.6 ค่าลักษณะเนื้อสัมผัสของชิฟฟอนเค้กเมื่อใช้แป้งฟักทองแทนที่แป้งสาลีที่ระดับแตกต่างกัน

ปริมาณแป้งฟักทองที่ใช้ แทนที่แป้งสาลี (%)	วัดค่า ซ้ำที่	Hardness (N)	Cohesiveness	Springiness (mm)	Chewiness (Nm)
0	1	18.56	0.477	10.46	0.113
	2	18.44	0.489	10.41	0.123
	3	18.23	0.456	10.35	0.114
	4	18.64	0.447	10.43	0.123
	5	18.55	0.465	10.58	0.145
10	1	20.76	0.453	9.56	0.126
	2	20.46	0.441	9.65	0.134
	3	20.24	0.457	9.37	0.123
	4	21.03	0.457	9.87	0.134
	5	21.05	0.461	9.36	0.135
20	1	22.45	0.401	9.11	0.078
	2	22.08	0.412	9.01	0.074
	3	22.46	0.416	9.16	0.077
	4	22.77	0.411	9.14	0.081
	5	22.78	0.418	9.23	0.083

จากข้อมูลดังกล่าว พิจารณารายละเอียดของข้อมูลได้ ดังนี้

- (1) ตัวแปรต้นหรือปัจจัยที่นิสิตปริญญาตรีสนใจ มี 1 ปัจจัย คือ ระดับแบ่งฟักทองที่ใช้แทนที่แบ่งสาลิ ดังนั้น จึงเป็นการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA)
- (2) ระดับของตัวแปรต้น เท่ากับ 3 ระดับ (0, 10 และ 20%) ดังนั้น จึงมี 3 สิ่งทดลอง และระดับแบ่งฟักทอง 0% คือ สิ่งทดลองควบคุม
- (3) ตัวแปรตาม มี 4 ตัว คือ ค่า Hardness, ค่า Cohesiveness, ค่า Springiness และค่า Chewiness
- (4) ต้องวิเคราะห์ One-Way ANOVA จำนวน 4 ครั้ง โดยวิเคราะห์ตัวแปรตามทีละตัว

8.6.2 การเขียนสมมติฐานสำหรับการวิเคราะห์ความแปรปรวนแบบทางเดียว

สามารถเขียนสมมติฐานสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ 4 ข้อตามจำนวนตัวแปรตามที่ต้องการวิเคราะห์ ดังนี้

สมมติฐานข้อที่ 1 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย Hardness

- | | | |
|-------|---|---|
| H_0 | : | ค่าเฉลี่ย Hardness ของชิฟฟอนเค้กทั้ง 3 สิ่งทดลองไม่แตกต่างกัน |
| H_1 | : | ค่าเฉลี่ย Hardness ของชิฟฟอนเค้กอย่างน้อย 2 สิ่งทดลองแตกต่างกัน |

สมมติฐานข้อที่ 2 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย Cohesiveness

- | | | |
|-------|---|---|
| H_0 | : | ค่าเฉลี่ย Cohesiveness ของชิฟฟอนเค้กทั้ง 3 สิ่งทดลองไม่แตกต่างกัน |
| H_1 | : | ค่าเฉลี่ย Cohesiveness ของชิฟฟอนเค้กอย่างน้อย 2 สิ่งทดลองแตกต่างกัน |

สมมติฐานข้อที่ 3 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย Springiness

- | | | |
|-------|---|--|
| H_0 | : | ค่าเฉลี่ย Springiness ของชิฟฟอนเค้กทั้ง 3 สิ่งทดลองไม่แตกต่างกัน |
| H_1 | : | ค่าเฉลี่ย Springiness ของชิฟฟอนเค้กอย่างน้อย 2 สิ่งทดลองแตกต่างกัน |

สมมติฐานข้อที่ 4 สำหรับทดสอบความแตกต่างของค่าเฉลี่ย Chewiness

- | | | |
|-------|---|--|
| H_0 | : | ค่าเฉลี่ย Chewiness ของชิฟฟอนเค้กทั้ง 3 สิ่งทดลองไม่แตกต่างกัน |
| H_1 | : | ค่าเฉลี่ย Chewiness ของชิฟฟอนเค้กอย่างน้อย 2 สิ่งทดลองแตกต่างกัน |

8.6.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ One-Way ANOVA

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลปริมาณแป้งฟักทองและค่าลักษณะเนื้อสัมผัสของชิฟฟอนเค้กแต่ละสิ่งทดลอง ให้กำหนดรายละเอียดตัวแปรในหน้าต่าง Variable View ดังภาพที่ 8.10 ในกรณีนี้ใช้ตัวแปร var จำนวน 5 ตัวแปรสำหรับป้อนข้อมูล 5 ตัวได้แก่ ปริมาณแป้งฟักทอง, ค่า Hardness, ค่า Cohesiveness, ค่า Springiness และค่า Chewiness ดังนั้นกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ปริมาณแป้งฟักทอง	Flour	Pumpkin Flour	0	ไม่ต้องกำหนด เพราะใช้ข้อมูลจริงในการป้อน
ค่า Hardness	Hardness	ไม่ต้องกำหนด	2	
ค่า Cohesiveness	Cohesiveness	ไม่ต้องกำหนด	3	
ค่า Springiness	Springiness	ไม่ต้องกำหนด	2	
ค่า Chewiness	Chewiness	ไม่ต้องกำหนด	3	

หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้

	Name	Type	Width	Decimals	Label	Values	Missing	Column
1	Flour	Numeric	8	0	Pumpkin Flour	None	None	8
2	Hardness	Numeric	8	2		None	None	8
3	Cohesiveness	Numeric	8	3		None	None	8
4	Springiness	Numeric	8	2		None	None	8
5	Chewiness	Numeric	8	3		None	None	8
6								
7								
8								
9								
10								
11								
12								
13								

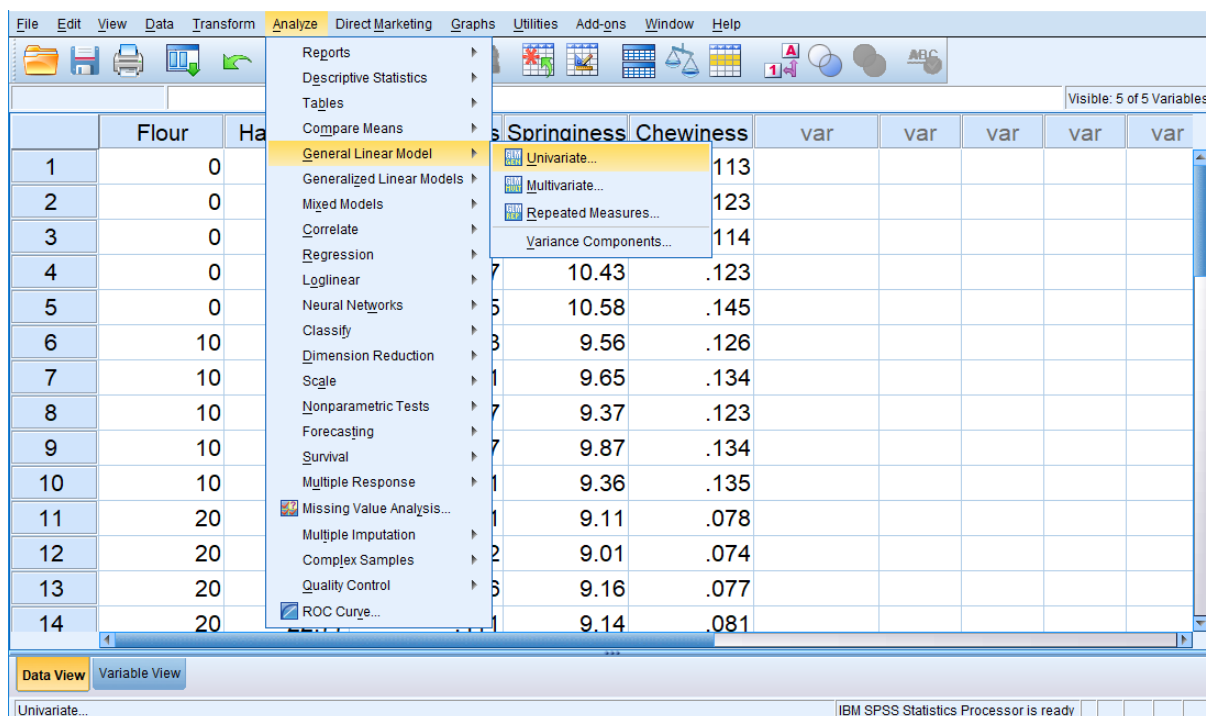
ภาพที่ 8.10 กำหนดรายละเอียดตัวแปรในหน้าต่าง Variable View

(2) ป้อนข้อมูลปริมาณแป้งฟักทอง, ค่า Hardness, ค่า Cohesiveness, ค่า Springiness และค่า Chewiness ในหน้าต่าง Data View ดังภาพที่ 8.11 โดยป้อนข้อมูลเรียงตามปริมาณแป้งฟักทองที่ใช้

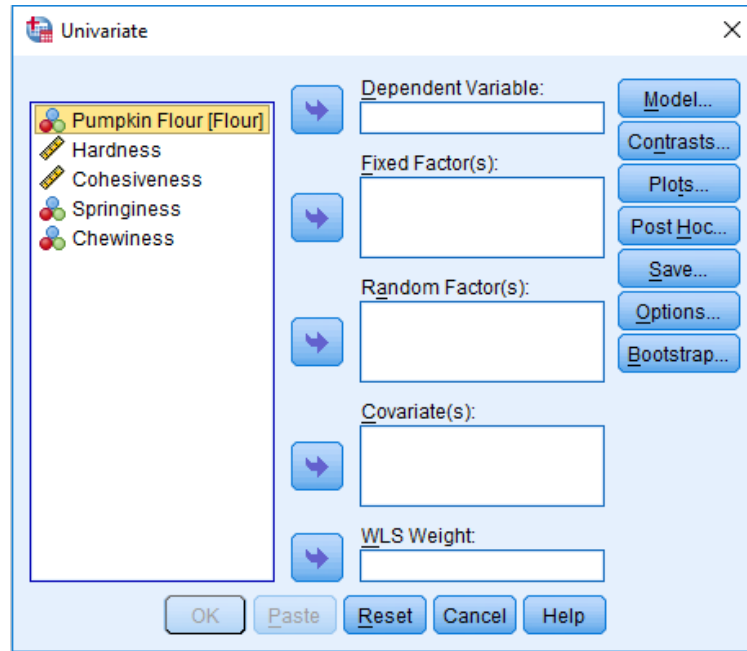
	Flour	Hardness	Cohesiveness	Springiness	Chewiness	var	var	var	var	var
1	0	18.56	.477	10.46	.113					
2	0	18.44	.489	10.41	.123					
3	0	18.23	.456	10.35	.114					
4	0	18.64	.447	10.43	.123					
5	0	18.55	.465	10.58	.145					
6	10	20.76	.453	9.56	.126					
7	10	20.46	.441	9.65	.134					
8	10	20.24	.457	9.37	.123					
9	10	21.03	.457	9.87	.134					
10	10	21.05	.461	9.36	.135					
11	20	22.45	.401	9.11	.078					
12	20	22.08	.412	9.01	.074					
13	20	22.46	.416	9.16	.077					
14	20	22.77	.411	9.14	.081					

ภาพที่ 8.11 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze ► General Linear Model ► Univariate...** ดังภาพที่ 8.12 จะปรากฏหน้าต่างดังภาพที่ 8.13 จะเห็นว่าในกล่องซ้ายมือจะแสดงชื่อของตัวแปรแต่ละตัว

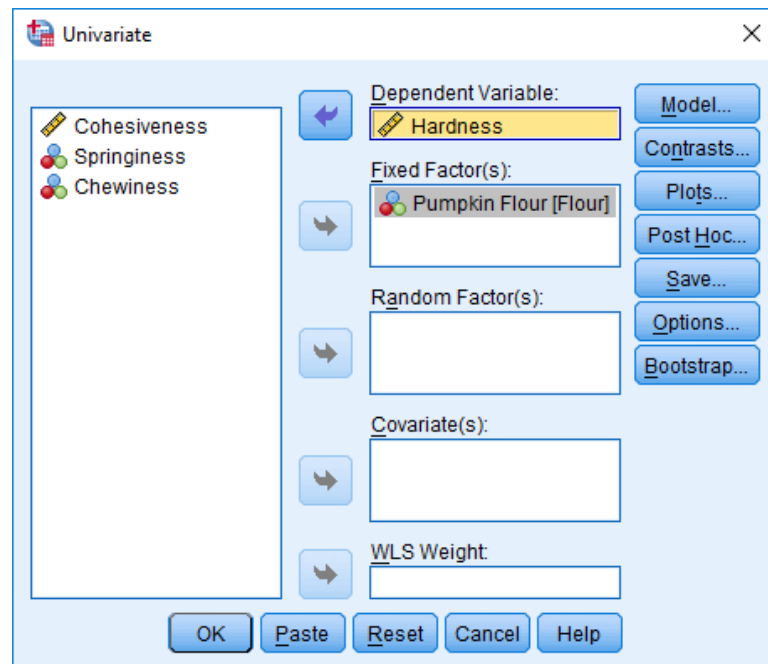


ภาพที่ 8.12 เมนูหน้าจอคำสั่ง Analyze ► General Linear Model ► Univariate...



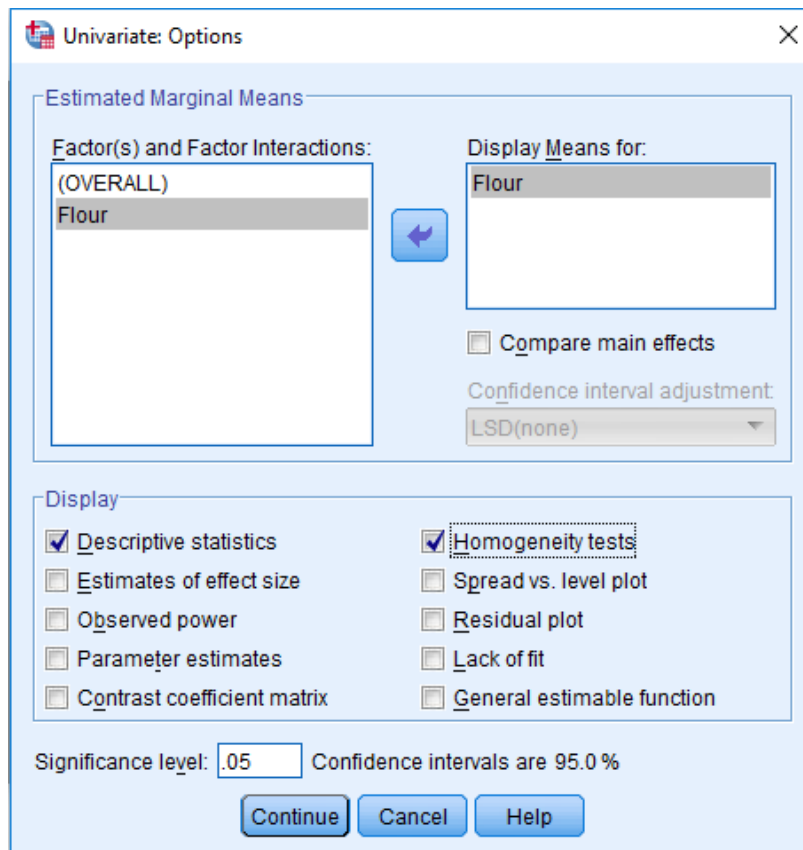
ภาพที่ 8.13 หน้าจอคำสั่ง Univariate

(4) เลือก ตัวแปร Pumpkin Flour [Flour] (ตัวแปรต้น) ใส่ในกล่อง Fixed Factor(s) และเลือกตัวแปรตาม ได้แก่ ตัวแปร Hardness (ตัวแปรตาม) ใส่ในกล่อง Dependent Variable จะได้ดังภาพที่ 8.14 (ในกรณีนี้มีตัวแปรตาม 4 ตัว จะต้องวิเคราะห์ความแปรปรวน 4 ครั้ง ครั้งละ 1 ตัวแปรตาม)



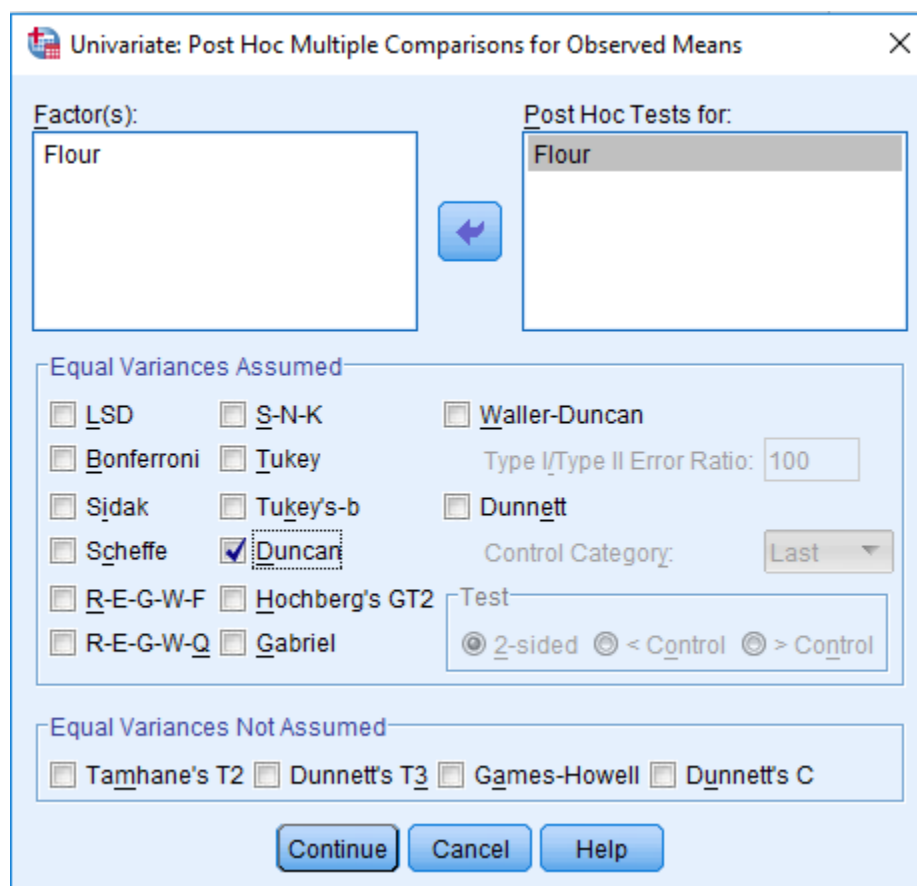
ภาพที่ 8.14 เลือกตัวแปร Hardness ใส่ในกล่อง Dependent Variable และเลือกตัวแปร Pumpkin Flour [Flour] ใส่ในกล่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Options...** จะได้หน้าจอตั้งภาพที่ 8.15 ให้กดเลือก **ตัวแปร Flour** จากกล่องซ้ายมือไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 8.15 หน้าจอคำสั่ง Univariate : Options

(6) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 8.16 ให้กดเลือก **ตัวแปร Flour** ในกล่องซ้ายมือ (Factor(s)) ใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธีอื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 8.17



ภาพที่ 8.16 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

(7) ทำการวิเคราะห์ความแปรปรวน ค่า Cohesiveness, ค่า Springiness และค่า Chewiness ตามขั้นตอนที่ (1) – (6) จะได้ผลลัพธ์แสดงดังภาพที่ 8.18, 8.19 และ 8.20 ตามลำดับ

Univariate Analysis of Variance

Between-Subjects Factors		
Pumpkin Flour		N
0		5
10		5
20		5

Descriptive Statistics

Dependent Variable: Hardness

Pumpkin Flour	Mean	Std. Deviation	N
0	18.4840	.15884	5
10	20.7080	.35492	5
20	22.5080	.28787	5
Total	20.5667	1.72311	15

Levene's Test of Equality of Error Variances^a

Dependent Variable: Hardness

F	df1	df2	Sig.
1.858	2	12	.198

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Flour

Tests of Between-Subjects Effects

Dependent Variable: Hardness

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	40.631 ^a	2	20.316	260.379	.000
Intercept	6344.817	1	6344.817	81319.477	.000
Flour	40.631	2	20.316	260.379	.000
Error	.936	12	.078		
Total	6386.384	15			
Corrected Total	41.568	14			

a. R Squared = .977 (Adjusted R Squared = .974)

Estimated Marginal Means

Pumpkin Flour

Dependent Variable: Hardness

Pumpkin Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	18.484	.125	18.212	18.756
10	20.708	.125	20.436	20.980
20	22.508	.125	22.236	22.780

Post Hoc Tests

Pumpkin Flour

Homogeneous Subsets

Hardness

Duncan^{a, b}

Pumpkin Flour	N	Subset		
		1	2	3
0	5	18.4840		
10	5		20.7080	
20	5			22.5080
Sig.		1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .078.

a. Uses Harmonic Mean Sample Size = 5.000.

b. Alpha = .05.

ภาพที่ 8.17 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่า Hardness ด้วยคำสั่ง Univariate

จากภาพที่ 8.17 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 3 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.198 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่า Hardness ของทั้ง 3 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Flour มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ย Hardness ของซีฟอนเค้กอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 3 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a, b และ c

Univariate Analysis of Variance

Between-Subjects Factors

Pumpkin Flour	N
0	5
10	5
20	5

Descriptive Statistics

Dependent Variable: Cohesiveness

Pumpkin Flour	Mean	Std. Deviation	N
0	.46680	.016649	5
10	.45380	.007694	5
20	.41160	.006580	5
Total	.44407	.026521	15

Levene's Test of Equality of Error Variances^a

Dependent Variable: Cohesiveness

F	df1	df2	Sig.
3.004	2	12	.088

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Flour

Tests of Between-Subjects Effects

Dependent Variable: Cohesiveness

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	.008 ^a	2	.004	32.900	.000
Intercept	2.958	1	2.958	23370.514	.000
Flour	.008	2	.004	32.900	.000
Error	.002	12	.000		
Total	2.968	15			
Corrected Total	.010	14			

a. R Squared = .846 (Adjusted R Squared = .820)

Estimated Marginal Means

Pumpkin Flour

Dependent Variable: Cohesiveness

Pumpkin Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	.467	.005	.456	.478
10	.454	.005	.443	.465
20	.412	.005	.401	.423

Post Hoc Tests

Pumpkin Flour

Homogeneous Subsets

Cohesiveness

Duncan^{a,b}

Pumpkin Flour	N	Subset	
		1	2
20	5	.41160	
10	5		.45380
0	5		.46680
Sig.		1.000	.093

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .000.

a. Uses Harmonic Mean Sample Size = 5.000.
b. Alpha = .05.

ภาพที่ 8.18 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่า Cohesiveness ด้วยคำสั่ง Univariate

จากภาพที่ 8.18 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 3 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.088 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่า Cohesiveness ของทั้ง 3 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Flour มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ย Cohesiveness ของชิฟฟอนเค้กอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 2 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a และ b

Univariate Analysis of Variance

Between-Subjects Factors

	N
Pumpkin Flour 0	5
10	5
20	5

Descriptive Statistics

Dependent Variable: Springiness

Pumpkin Flour	Mean	Std. Deviation	N
0	10.4460	.08503	5
10	9.5620	.21230	5
20	9.1300	.08031	5
Total	9.7127	.58156	15

Levene's Test of Equality of Error Variances^a

Dependent Variable: Springiness

F	df1	df2	Sig.
2.664	2	12	.110

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Flour

Tests of Between-Subjects Effects

Dependent Variable: Springiness

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	4.500 ^a	2	2.250	114.891	.000
Intercept	1415.038	1	1415.038	72257.280	.000
Flour	4.500	2	2.250	114.891	.000
Error	.235	12	.020		
Total	1419.773	15			
Corrected Total	4.735	14			

a. R Squared = .950 (Adjusted R Squared = .942)

Estimated Marginal Means

Pumpkin Flour

Dependent Variable: Springiness

Pumpkin Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	10.446	.063	10.310	10.582
10	9.562	.063	9.426	9.698
20	9.130	.063	8.994	9.266

Post Hoc Tests

Pumpkin Flour

Homogeneous Subsets

Springiness

Duncan^{a, b}

Pumpkin Flour	N	Subset		
		1	2	3
20	5	9.1300		
10	5		9.5620	
0	5			10.4460
Sig.		1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.

Based on observed means.

The error term is Mean Square(Error) = .020.

a. Uses Harmonic Mean Sample Size = 5.000.

b. Alpha = .05.

ภาพที่ 8.19 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่า Springiness ด้วยคำสั่ง Univariate

จากภาพที่ 8.19 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 3 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.110 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่า Springiness ของทั้ง 3 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Flour มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ย Springiness ของซีฟอนเค้ก อย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 3 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a, b และ c

Univariate Analysis of Variance

Between-Subjects Factors		
Pumpkin Flour		N
0		5
10		5
20		5

Descriptive Statistics

Dependent Variable: Chewiness

Pumpkin Flour	Mean	Std. Deviation	N
0	.12360	.012876	5
10	.13040	.005505	5
20	.07860	.003507	5
Total	.11087	.025011	15

Levene's Test of Equality of Error Variances^a

Dependent Variable: Chewiness

F	df1	df2	Sig.
1.658	2	12	.231

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Flour

Tests of Between-Subjects Effects

Dependent Variable: Chewiness

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	.008 ^a	2	.004	57.036	.000
Intercept	.184	1	.184	2654.097	.000
Flour	.008	2	.004	57.036	.000
Error	.001	12	6.947E-5		
Total	.193	15			
Corrected Total	.009	14			

a. R Squared = .905 (Adjusted R Squared = .889)

Estimated Marginal Means

Pumpkin Flour

Dependent Variable: Chewiness

Pumpkin Flour	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
0	.124	.004	.115	.132
10	.130	.004	.122	.139
20	.079	.004	.070	.087

Post Hoc Tests

Pumpkin Flour

Homogeneous Subsets

Chewiness

Duncan^{a, b}

Pumpkin Flour	N	Subset	
		1	2
20	5	.07860	
0	5		.12360
10	5		.13040
Sig.		1.000	.221

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = 6.95 E-005.

a. Uses Harmonic Mean Sample Size = 5.000.

b. Alpha = .05.

ภาพที่ 8.20 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนของค่า Chewiness ด้วยคำสั่ง Univariate

จากภาพที่ 8.20 ให้พิจารณาผลลัพธ์จากการวิเคราะห์ One-Way ANOVA 3 ขั้นตอน ดังนี้

(1) พิจารณาตาราง Levene's Test of Equality of Error Variances ที่ใช้ทดสอบการกระจายตัวของข้อมูลทั้ง 3 กลุ่ม พบว่าค่า Sig. มีค่าเท่ากับ 0.231 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทั้ง 3 กลุ่มมีการกระจายตัวไม่แตกต่างกันที่ระดับนัยสำคัญ 0.05

(2) พิจารณาตาราง Tests of Between-Subjects Effects ที่ใช้ทดสอบความแตกต่างของค่า Chewiness ของทั้ง 3 สิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Flour มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้นปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ย Chewiness ของซีฟอนเค้ก อย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05 ขั้นตอนถัดไปจึงต้องทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ

(3) พิจารณาตาราง Post Hoc Tests ที่ใช้สำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ พบว่า ความแตกต่างค่าเฉลี่ยของสิ่งทดลองแบ่งออกเป็น 2 Subset ดังนั้นตัวอักษรภาษาอังกฤษที่จะใช้แสดงกลุ่มที่แตกต่าง คือ a และ b

8.6.4 การนำเสนอผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ CRD

จากผลลัพธ์การวิเคราะห์ความแปรปรวน ค่า Hardness, ค่า Cohesiveness, ค่า Springiness และค่า Chewiness ดังแสดงในภาพที่ 8.17, 8.18, 8.19 และ 8.20 ตามลำดับนั้น นักวิจัยสามารถสรุปผลการวิเคราะห์ทั้ง 4 ค่าพารามิเตอร์ในรูปแบบตารางเพียง 1 ตาราง โดยในการเขียนผลการวิเคราะห์ของแต่ละพารามิเตอร์นั้นจะนำข้อมูลค่าเฉลี่ย (Mean) และค่าส่วนเบี่ยงเบนมาตรฐาน (Std. Deviation) ของแต่ละสิ่งทดลองจากข้อมูลในตาราง Descriptive Statistics และผลการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณจากข้อมูลในตาราง Post Hoc Tests มาใช้ในการเขียนผล ดังแสดงในตารางที่ 8.7

ตารางที่ 8.7 ค่าลักษณะเนื้อสัมผัสของชิฟฟอนเค้กเมื่อใช้แป้งฟักทองแทนที่แป้งสาลีที่ระดับแตกต่างกัน

ปริมาณแป้งฟักทองที่ใช้แทนที่แป้งสาลี (%)	Hardness (N)	Cohesiveness	Springiness (mm)	Chewiness (Nm)
0	18.48 ± 0.16c	0.467 ± 0.017a	10.45 ± 0.09a	0.12 ± 0.01a
10	20.71 ± 0.35b	0.454 ± 0.008a	9.56 ± 0.21b	0.13 ± 0.01a
20	22.51 ± 0.29a	0.412 ± 0.01b	9.13 ± 0.08c	0.08 ± 0.00b

^{a-c} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

ข้อสังเกต ในการเรียงลำดับการใช้ตัวอักษรภาษาอังกฤษเพื่อแสดงกลุ่มที่แตกต่างของสิ่งทดลองนั้น นักวิจัยสามารถเรียงลำดับ a, b, c, ... จากค่าเฉลี่ยมากไปน้อย หรือ จากค่าเฉลี่ยน้อยไปมาก ทั้งนี้ นักวิจัยต้องเลือกอย่างใดอย่างหนึ่งและต้องใช้เกณฑ์นี้กับทุกค่าตัวแปรตามในงานวิจัยเรื่องนั้นเสมอ



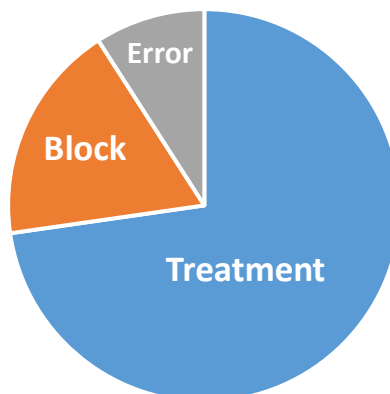
บทที่ 9

การวางแผนการทดลอง แบบสุ่มในบล็อกสมบูรณ์

การวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (Randomized Complete Block Design : RCBD หรือ RCB หรือ RBD) นิยมนำมาใช้ในกรณีที่หน่วยทดลอง (Experimental unit) ไม่มีความสม่ำเสมอ เกิดความผันแปรในหน่วยทดลองก่อนที่จะกำหนดสิ่งทดลองหรือทรีทเมนต์ (Treatment) ให้แก่หน่วยทดลอง ดังนั้นเพื่อให้ความผันแปรที่เกิดขึ้นเป็นผลเนื่องมาจากสิ่งทดลองเพียงอย่างเดียว จึงทำการควบคุมความผันแปรทางด้านหน่วยทดลองโดยการแบ่งเป็นบล็อก (Block)

วัตถุประสงค์ของการแบ่งหน่วยทดลองเป็นบล็อก คือ เพื่อให้หน่วยทดลองภายในบล็อกเดียวกันมีความแตกต่างกันน้อยที่สุด และหน่วยทดลองที่อยู่ต่างบล็อกกันให้มีความแตกต่างกันมากที่สุด นั่นคือให้หน่วยทดลองภายในบล็อกเดียวกันมีความสม่ำเสมอมากที่สุดก่อนที่จะกำหนดสิ่งทดลองให้แก่หน่วยทดลอง การวางแผนการทดลองแบบ RCBD นี้แต่ละสิ่งทดลองจะมีจำนวนซ้ำ (Replication) เท่ากัน ทั้งนี้ในแต่ละบล็อกจะมีสิ่งทดลองครบทุกสิ่งทดลอง

การวางแผนการทดลองแบบ RCBD มีการผันแปรของ 2 ปัจจัย โดยการทดลองนั้นมักมีมากกว่า 2 สิ่งทดลองขึ้นไป แหล่งของความแปรปรวนในการวางแผนการทดลองแบบ RCBD มี 3 แหล่ง คือ จากสิ่งทดลอง (Treatment), จากบล็อก (Block) และ จากความคลาดเคลื่อนจากการทดลอง (Experimental error) ดังภาพที่ 9.1 เงื่อนไขของการวางแผนการทดลองแบบ RCBD คือ ต้องไม่มีปฏิสัมพันธ์ (Interaction) ระหว่างบล็อกและสิ่งทดลอง



ภาพที่ 9.1 การแบ่งความแปรปรวนของแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD)

9.1 ข้อดีและข้อเสียของการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD)

ข้อดีของการวางแผนการทดลองแบบ RCBD

- สามารถจัดหน่วยทดลองให้เป็นบล็อกได้ค่อนข้างง่าย
- ทำการวิเคราะห์ข้อมูลง่าย
- ถ้ามีข้อมูลสูญหาย สามารถประมาณค่าข้อมูลสูญหายได้

ข้อเสียของการวางแผนการทดลองแบบ RCBD

- เมื่อมีความผันแปรระหว่างหน่วยทดลองภายในบล็อกเดียวกันสูง จะทำให้ความคลาดเคลื่อนในการทดลองสูง ซึ่งมักเกิดกับการทดลองที่มีจำนวนสิ่งทดลองมากเกินไป
- ถ้าหน่วยทดลองมีความสม่ำเสมอ การวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) จะมีประสิทธิภาพมากกว่าการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD)

9.2 การบล็อก (Blocking)

การบล็อก คือ การจัดหน่วยทดลองที่มีความคล้ายคลึงกันหรือเหมือนกันให้อยู่ในบล็อกเดียวกันหรือกลุ่มเดียวกันเพื่อให้มีความแปรปรวนภายในบล็อกเดียวกันน้อยที่สุด และหน่วยทดลองที่อยู่ต่างบล็อกกันมีความแตกต่างกันมากที่สุดเพื่อให้มีความแปรปรวนระหว่างบล็อกมากที่สุด ในการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) บล็อก คือ ขั้วของการทดลอง เช่น การบล็อกแหล่งซื้อเมล็ดถั่วเหลือง จะได้เมล็ดถั่วเหลืองที่ซื้อจากตลาด A อยู่ในบล็อกที่หนึ่ง และ เมล็ดถั่วเหลืองที่ซื้อจากตลาด B อยู่ในบล็อกที่สอง

Block ที่ 1 คือ เมล็ดถั่วเหลืองที่ซื้อจากตลาด A

Block ที่ 2 คือ เมล็ดถั่วเหลืองที่ซื้อจากตลาด B

ประโยชน์ของการทำบล็อก

- (1) ช่วยเพิ่มความถูกต้องให้แก่การทดลอง เนื่องจากความผันแปรระหว่างบล็อกจะถูกดึงออกจากความคลาดเคลื่อนของการทดลอง จึงเหลือเฉพาะความคลาดเคลื่อนจากการทดลองเท่านั้น
- (2) เป็นการเปรียบเทียบอิทธิพลของสิ่งทดลองที่มีต่อหน่วยทดลองที่มีลักษณะเหมือนกัน ซึ่งจัดให้หน่วยทดลองที่เหมือนกันอยู่ในบล็อกเดียวกัน ส่วนหน่วยทดลองที่ต่างกันอย่างอยู่ต่างบล็อกกัน
- (3) ช่วยขยายขอบเขตของการทดลอง เช่น สถานที่ทำการทดลองอยู่คนละจังหวัดกัน ทำการทดลองในเวลาที่แตกต่างกัน ในกรณีนี้ สถานที่และเวลา เป็น บล็อก

9.3 ตัวแบบ (Model)

ในการวางแผนการทดลองแบบ RCBD การวิเคราะห์ความแปรปรวนสำหรับตัวแบบ (Model) เป็นการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) เนื่องจากสามารถแยกอิทธิพลของบล็อกออกมาเป็นแหล่งความแปรปรวนหนึ่งนอกจากแหล่งความแปรปรวนอันเนื่องมาจากอิทธิพลของสิ่งทดลอง แต่ยังคงเป็นการวิเคราะห์ความแปรปรวนที่มีปัจจัยเดียว (Single factor analysis of variance) เนื่องจากนักวิจัยสนใจเฉพาะความแตกต่างระหว่างสิ่งทดลอง ทั้งนี้สาเหตุที่ต้องจัดเป็นบล็อกเนื่องจากหน่วยทดลองไม่มีความสม่ำเสมอ ตัวแบบ (Model) ของการวางแผนการทดลองแบบ RCBD เขียนได้ดังนี้

$$y_{ij} = \mu + \tau_i + \beta_j + \varepsilon_{ij} \quad ; \quad \begin{array}{l} i = 1, 2, 3, \dots, a \\ j = 1, 2, 3, \dots, b \end{array}$$

โดยที่	y_{ij}	คือ	ค่าสังเกตที่วัดได้จากสิ่งทดลองที่ i บล็อกที่ j
	μ	คือ	ค่าเฉลี่ยข้อมูลทั้งหมด (Overall mean)
	τ_i	คือ	อิทธิพล (Effect) จากสิ่งทดลองที่ i
	β_j	คือ	อิทธิพล (Effect) จากบล็อกที่ j
	ε_{ij}	คือ	ค่าความคลาดเคลื่อนจากการทดลอง (Experimental error หรือ Random error) ที่อยู่ในค่าสังเกตของสิ่งทดลองที่ i บล็อกที่ j

ในงานวิจัยและพัฒนาผลิตภัณฑ์นิยมใช้แผนการทดลองแบบ RCBD ในการประเมินคุณภาพทางด้านประสาทสัมผัส เช่น ในขั้นตอนการคัดเลือกสูตรพื้นฐาน การพัฒนาสูตรและกรรมวิธีการผลิต การศึกษาอายุการเก็บรักษาของผลิตภัณฑ์ โดยบล็อกอาจจะเป็นผู้ชิม การทำซ้ำ หรือ สถานที่ทดสอบก็ได้ เช่น ในการประเมินคุณภาพทางด้านประสาทสัมผัสตัวอย่างเค้กที่ลดปริมาณไขมัน 3 ระดับ (0, 10 และ 20%) นักวิจัยทำการทดสอบ 1 ครั้งและใช้ผู้ทดสอบชิมจำนวน 50 คน ในกรณีนี้จะถือว่าสิ่งทดลอง (Treatment) มี 3 สิ่งทดลอง คือ เค้กที่ลดปริมาณไขมันที่ 0, 10 และ 20% และบล็อก (Block) คือ ผู้ทดสอบชิม โดยกรณีนี้บล็อก แทนการทำซ้ำ (Replication) ซึ่งในหนังสือตำราสถิติบางเล่มจะถือว่าเป็นการวัดค่าซ้ำโดยจะมีการวิเคราะห์ความแปรปรวนเหมือนกันทุกประการ

หากนำตัวอย่างชุดเดียวที่ได้จากการเตรียมครั้งเดียวมาประเมินคุณภาพทางประสาทสัมผัส จะไม่สามารถประมาณค่าความคลาดเคลื่อนของการทดลอง (Experimental error) ได้ แต่จะได้ความคลาดเคลื่อนของการวัดค่า (Measurement error) แทน ซึ่งผลที่ได้จากการวิเคราะห์จะสามารถสรุปครอบคลุมได้เฉพาะ

ตัวอย่างนั้นๆ หรือเฉพาะการทดลองครั้งนั้น ไม่สามารถสรุปให้ครอบคลุมตัวผลิตภัณฑ์ได้จนกว่าจะได้มีการประเมินผลิตภัณฑ์ต่างชุดกัน (เตรียมตัวอย่างใหม่เพื่อทำซ้ำ)

ตารางที่ 9.1 ผังแสดงค่าสังเกตของตัวอย่างสำหรับการวิเคราะห์ความแปรปรวนของข้อมูลที่ได้จากการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์

สิ่งทดลองที่	บล็อก						ผลรวม
	1	2	...	j	...	b	
1	y_{11}	y_{12}	...	y_{1j}	...	y_{1b}	$y_{1\cdot}$
2	y_{21}	y_{22}	...	y_{2j}	...	y_{2b}	$y_{2\cdot}$
.	.	.	...	.	...	.	.
.	.	.	...	.	...	.	.
i	y_{i1}	y_{i2}	...	y_{ij}	...	y_{ib}	$y_{i\cdot}$
.	.	.	...	.	...	.	.
.	.	.	...	.	...	.	.
a	y_{a1}	y_{a2}	...	y_{aj}	...	y_{ab}	$y_{a\cdot}$
ผลรวม	$y_{\cdot 1}$	$y_{\cdot 2}$	...	$y_{\cdot j}$	...	$y_{\cdot b}$	$y_{\cdot\cdot}$

โดยที่	y_{ij}	คือ	ค่าสังเกตที่วัดได้จากสิ่งทดลองที่ i บล็อกที่ j
	$y_{i\cdot}$	คือ	ผลรวมของค่าสังเกตทั้งหมดที่ได้รับสิ่งทดลองที่ i
	$\overline{y_{i\cdot}}$	คือ	ค่าเฉลี่ยของค่าสังเกตทั้งหมดที่ได้รับสิ่งทดลองที่ i
	$y_{\cdot j}$	คือ	ผลรวมของค่าสังเกตทั้งหมดในบล็อกที่ j
	$\overline{y_{\cdot j}}$	คือ	ค่าเฉลี่ยของค่าสังเกตทั้งหมดในบล็อกที่ j
	$y_{\cdot\cdot}$	คือ	ผลรวมของค่าสังเกตทั้งหมด
	$\overline{y_{\cdot\cdot}}$	คือ	ค่าเฉลี่ยของค่าสังเกตทั้งหมด
	$N = ab$	คือ	ผลรวมของจำนวนของค่าสังเกตทั้งหมด

9.4 การตั้งสมมติฐานเพื่อการทดสอบความสัมพันธ์หรือทดสอบเปรียบเทียบค่าเฉลี่ย

วัตถุประสงค์ที่แท้จริงของการทดสอบสมมติฐาน คือ การทดสอบเกี่ยวกับอิทธิพลของสิ่งทดลองที่มีต่อค่าสังเกตมากกว่าที่จะทดสอบอิทธิพลของบล็อก ทั้งนี้หากนักวิจัยต้องการทดสอบอิทธิพลของบล็อกที่มีต่อค่าสังเกตก็สามารถทำได้ ข้อควรสังเกต คือ ถ้าค่าสถิติ F ในตาราง ANOVA ที่ใช้ทดสอบอิทธิพลของบล็อกมีนัยสำคัญ (มีค่าน้อยกว่าค่า α ที่นักวิจัยกำหนดไว้) แสดงว่าการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) จะช่วยเพิ่มประสิทธิภาพได้ดีกว่าการวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD)

ในหนังสือเล่มนี้จะอธิบายเฉพาะวิธีการทดสอบเกี่ยวกับสิ่งทดลองและบล็อกเมื่อเป็นอิทธิพลแบบคงที่ (Fixed effect model) เท่านั้น ซึ่งสมมติฐานงานวิจัยที่ต้องการทดสอบในตัวแบบอิทธิพลคงที่เขียนได้ ดังนี้

(1) สมมติฐานสำหรับทดสอบอิทธิพลของสิ่งทดลอง

H_0 : ค่าเฉลี่ยของตัวอย่างทุกสิ่งทดลองไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยของตัวอย่างอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\mu_{.1} = \mu_{.2} = \dots = \mu_{.a}$

H_1 : $\mu_{.i} \neq \mu_{.j}$ อย่างน้อย 1 คู่; $i \neq j$

ถ้ายอมรับสมมติฐานหลัก (Accept H_0) แสดงว่าสิ่งทดลองมีค่าเฉลี่ย μ ไม่แตกต่างกัน หรือ สิ่งทดลองไม่มีอิทธิพลต่อค่าสังเกต และ ถ้าปฏิเสธสมมติฐานหลัก (Reject H_0) แสดงว่าสิ่งทดลองมีค่าเฉลี่ย μ แตกต่างกัน หรือ สิ่งทดลองมีอิทธิพลต่อค่าสังเกต ต้องทำการเปรียบเทียบความแตกต่างของค่าเฉลี่ยโดยใช้วิธีการทดสอบพหุคูณ (Multiple comparison tests) หรือ Post hoc tests กรณีที่ตัวแปรต้นหรือระดับของปัจจัยที่สนใจมี k กลุ่ม จะต้องทดสอบทั้งหมด kC_2 คู่ เช่น มี 3 อาชีพ ($k=3$) จะมี 3C_2 นั่นเอง

(2) สมมติฐานสำหรับทดสอบอิทธิพลของบล็อก

H_0 : ค่าเฉลี่ยของตัวอย่างแต่ละบล็อกไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยของตัวอย่างอย่างน้อย 2 บล็อกแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\mu_{.1} = \mu_{.2} = \dots = \mu_{.b}$

H_1 : $\mu_{.i} \neq \mu_{.j}$ อย่างน้อย 1 คู่; $i \neq j$

การวิเคราะห์ความแปรปรวนสำหรับข้อมูลที่ได้จากการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) เป็นการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) ผลการวิเคราะห์ความแปรปรวนสามารถสรุปเป็นตารางดังแสดงในตารางที่ 9.2

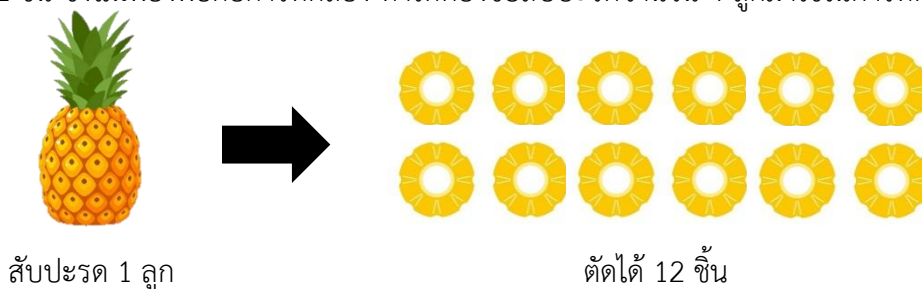
ตารางที่ 9.2 การวิเคราะห์ความแปรปรวนสำหรับการวางแผนการทดลองแบบ RCBD

แหล่งความแปรปรวน (Source of Variation)	องศาความเป็นอิสระ (Degree of Freedom)	ผลบวกกำลังสอง (Sum of Squares)	ค่ากำลังสองเฉลี่ย (Mean Squares)	F
ระหว่างสิ่งทดลอง (Between Treatment)	$a - 1$	$SSTr = \sum_{i=1}^a \frac{y_{i.}^2}{b} - \frac{y_{..}^2}{ab}$	$MSTr = \frac{SSTr}{a - 1}$	$\frac{MSTr}{MSE}$
ระหว่างบล็อก (Between Block)	$b - 1$	$SSB = \sum_{j=1}^b \frac{y_{.j}^2}{a} - \frac{y_{..}^2}{ab}$	$MSB = \frac{SSB}{b - 1}$	
ภายในสิ่งทดลองหรือความคลาดเคลื่อน (Within Treatment or Error)	$(a - 1)(b - 1)$	$SSE = SST - SSTr - SSB$	$MSE = \frac{SSE}{(a - 1)(b - 1)}$	
ยอดรวม (Total)	$ab - 1$	$SST = \sum_{i=1}^a \sum_{j=1}^b y_{ij}^2 - \frac{y_{..}^2}{ab}$		

a = จำนวนสิ่งทดลองหรือระดับของปัจจัย, b = จำนวนบล็อก

9.5 ตัวอย่างรูปแบบข้อมูลจากการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์

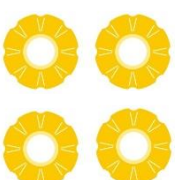
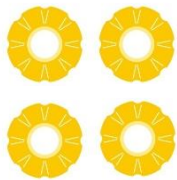
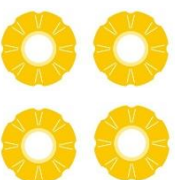
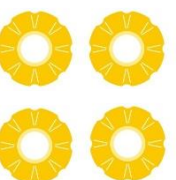
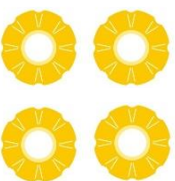
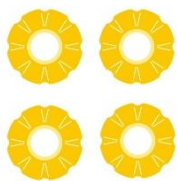
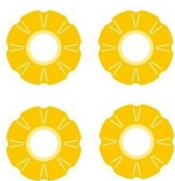
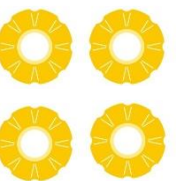
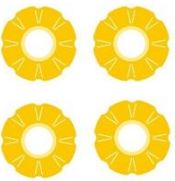
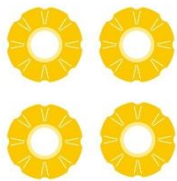
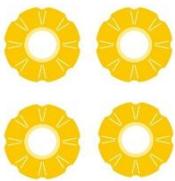
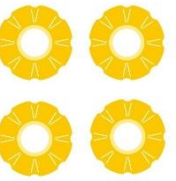
นักวิจัยต้องการศึกษาผลของอุณหภูมิในการทอดแบบสุญญากาศต่อคุณภาพของสับปะรดทอดในด้านปริมาณความชื้น, ปริมาณไขมัน และ ความกรอบ โดยศึกษาอุณหภูมิในการทอด 3 ระดับ คือ 100, 110 และ 120 องศาเซลเซียส ทั้งนี้จากการศึกษาเบื้องต้นนักวิจัยพบว่าแต่ละสิ่งทดลองจะต้องใช้สับปะรดทอดจำนวน 16 ชิ้นที่มีขนาดเท่ากันสำหรับวัดค่าคุณภาพต่างๆ ได้แก่ วิเคราะห์ปริมาณความชื้น 4 ชิ้น, วิเคราะห์ปริมาณไขมัน 4 ชิ้น และวัดค่าความกรอบ 8 ชิ้น ซึ่งในขั้นตอนการเตรียมสับปะรดพบว่า สับปะรด 1 ลูกสามารถตัดแต่งเป็นชิ้นที่เท่ากันได้ 12 ชิ้น ซึ่งไม่เพียงพอต่อการทดลอง ทำให้ต้องซื้อสับปะรดจำนวน 4 ลูกมาใช้ในการทดลอง



ทั้งนี้เมื่อนักวิจัยจะซื้อสับปะรดสายพันธุ์เดียวกันมาจากร้านค้าเดียวกัน ขนาดและน้ำหนักเท่ากัน แต่เป็นการยากที่จะควบคุมให้มีคุณภาพสม่ำเสมอได้ การทดลองดังกล่าวจึงจำเป็นต้องใช้การวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) กล่าวคือ ต้องมีการบล็อก (Block) ขึ้นสับปะรดเพื่อลดความคลาดเคลื่อนจากความไม่สม่ำเสมอดังกล่าว โดยจัดให้ทุกสิ่งทดลองมีโอกาสได้รับหน่วยทดลองที่เหมือนกันในแต่ละบล็อก

ในกรณีนี้ สิ่งทดลอง คือ อุณหภูมิในการทอด และ บล็อก คือ สับปะรดแต่ละลูก กล่าวคือ สับปะรด 1 ลูกเมื่อปอกเปลือกและตัดแต่งให้มีขนาดเท่ากันได้จำนวน 12 ชิ้นแล้วจะนำมาแบ่งออกเป็น 3 ส่วน เพื่อสุ่มไปทอดที่ 3 อุณหภูมิ (เท่ากับจำนวนสิ่งทดลอง) นั่นคือสับปะรด 1 ลูกจะได้รับสิ่งทดลองทั้ง 3 สิ่งทดลองโดยสุ่มและทำเช่นเดียวกันกับสับปะรดอีก 3 ลูก ซึ่งจะได้แผนผังการจัดวางสิ่งทดลองดังตารางที่ 9.3 จากตารางจะสังเกตได้ว่าแต่ละสิ่งทดลองจะได้สับปะรดจำนวน 16 ชิ้นซึ่งเพียงพอต่อการทดลองของนักวิจัย

ตารางที่ 9.3 ผังสำหรับการศึกษาอิทธิพลของอุณหภูมิในการทอดต่อคุณภาพสับปะรดทอดสุญญากาศ โดยใช้การวางแผนการทดลองแบบ RCBD

สิ่งทดลองที่ (อุณหภูมิในการทอด °C)	บล็อกที่ 1 	บล็อกที่ 2 	บล็อกที่ 3 	บล็อกที่ 4 
สิ่งทดลองที่ 1: อุณหภูมิ 100 °C				
สิ่งทดลองที่ 2 : อุณหภูมิ 110 °C				
สิ่งทดลองที่ 3 : อุณหภูมิ 120 °C				

จากตัวอย่างข้างต้น จึงสามารถสรุปรายละเอียดของข้อมูลได้ดังนี้

- (1) ปัจจัยที่ศึกษา มี 1 ปัจจัย คือ อุณหภูมิในการทอด
- (2) จำนวนสิ่งทดลอง (Treatment) เท่ากับ 3 สิ่งทดลอง (ตามระดับของอุณหภูมิในการทอด)
- (3) บล็อก (Block) คือ ลูกสับปะรด
- (4) จำนวนบล็อก เท่ากับ 4 บล็อก (เท่ากับจำนวนลูกสับปะรด)
- (5) ค่าสังเกต (ตัวแปรตาม) มี 3 ค่า คือ ปริมาณความชื้น, ปริมาณไขมัน และ ค่าความกรอบ
- (6) จำนวนหน่วยทดลอง (Experimental unit) ของแต่ละสิ่งทดลอง
 - ค่าความกรอบ มีจำนวนหน่วยทดลอง เท่ากับ 8 หน่วย (8 ชิ้น)
 - ปริมาณความชื้น มีจำนวนหน่วยทดลอง เท่ากับ 4 หน่วย (4 ชิ้น)
 - ปริมาณไขมัน มีจำนวนหน่วยทดลอง เท่ากับ 4 หน่วย (4 ชิ้น)

ดังนั้นในการสุ่มหน่วยทดลองเพื่อมาวัดค่าสังเกตทั้ง 3 ค่า ผู้วิจัยต้องสุ่มมาจากทุกบล็อกเพื่อให้หน่วยทดลองในแต่ละบล็อกมีความสม่ำเสมอ

9.6 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD

การวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) เป็นการวิเคราะห์ความแปรปรวนที่มีปัจจัยเดียว (Single factor analysis of variance) เนื่องจากนักวิจัยสนใจเฉพาะความแตกต่างระหว่างสิ่งทดลอง แต่ในการกำหนดตัวแบบ (Model) สำหรับการวิเคราะห์ความแปรปรวนต้องเป็นแบบสองทาง (Two-Way ANOVA) โดยความแปรปรวนที่อธิบายได้เกิดจากอิทธิพลหลัก (Main effect) 2 แหล่ง คือ จากสิ่งทดลอง (Treatment) และ จากบล็อก (Block) ทั้งนี้ในการวิเคราะห์จะไม่กำหนดอิทธิพลร่วม (Interaction) ระหว่างสิ่งทดลองและบล็อก เนื่องจากเป็นเงื่อนไขของการวางแผนการทดลองแบบ RCBD ดังนั้นคำสั่งที่ใช้วิเคราะห์ข้อมูล คือ

Analyze ➤ General Linear Model ➤ Univariate...

9.6.1 ตัวอย่างข้อมูล

นักวิจัยต้องการทราบว่าอุณหภูมิที่ใช้ในการทอด 3 ระดับ คือ 100, 110 และ 120 องศาเซลเซียส มีผลต่อปริมาณความชื้นของสับปะรดทอดหรือไม่ ทำการทดลองซ้ำ 3 ครั้งโดยแต่ละครั้งใช้สับปะรด 1 ลูก เนื่องจากสับปะรดแต่ละลูกมีคุณภาพไม่สม่ำเสมอ จึงวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) และในแต่ละครั้งของการทดลองจะสุ่มตัวอย่างสับปะรดแต่ละสิ่งทดลองมาวิเคราะห์ปริมาณความชื้น 2 ชิ้น ได้ข้อมูลดังตารางที่ 9.4

ตารางที่ 9.4 ปริมาณความชื้นของสับปะรดที่ทอดด้วยอุณหภูมิที่ต่างกัน

สับปะรดลูกที่ (Block)	การวัดค่าซ้ำที่	ปริมาณความชื้น (%) สับปะรดที่ทอดด้วยอุณหภูมิ		
		100 °C	110 °C	120 °C
1	1	8.5	7.5	5.1
	2	8.6	7.6	5.2
2	1	8.1	7.2	5.8
	2	8.2	7.1	5.7
3	1	9.4	7.8	5.4
	2	9.2	7.9	5.6

9.6.2 การเขียนสมมติฐานเพื่อทดสอบ

จากข้อมูลดังกล่าว สามารถสรุปรายละเอียดของข้อมูลได้ ดังนี้

- (1) ปัจจัยที่ศึกษา มี 1 ปัจจัย คือ อุณหภูมิทอด
- (2) จำนวนสิ่งทดลอง (Treatment) เท่ากับ 3 สิ่งทดลอง (ตามระดับอุณหภูมิทอด)
- (3) บล็อก (Block) คือ สับปะรดแต่ละลูก
- (4) จำนวนบล็อก เท่ากับ 3 บล็อก (เท่ากับจำนวนสับปะรด)
- (6) จำนวนหน่วยทดลอง (Experimental unit) ในแต่ละสิ่งทดลอง เท่ากับ 6 หน่วย
- (7) จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 18 หน่วย
- (8) ค่าสังเกต (ตัวแปรตาม) คือ ปริมาณความชื้น

นักวิจัยสามารถเขียนสมมติฐานทางสถิติสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ ดังนี้

(1) ทดสอบอิทธิพลของสิ่งทดลอง (อุณหภูมิที่ใช้ในการทอด) ที่มีต่อปริมาณความชื้น

H_0 : ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดทุกสิ่งทดลองไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดอย่างน้อย 2 สิ่งทดลองแตกต่างกัน

(2) ทดสอบอิทธิพลของบล็อก (สับปะรดแต่ละลูก) ที่มีต่อปริมาณความชื้น

H_0 : ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดทุกบล็อกไม่แตกต่างกัน

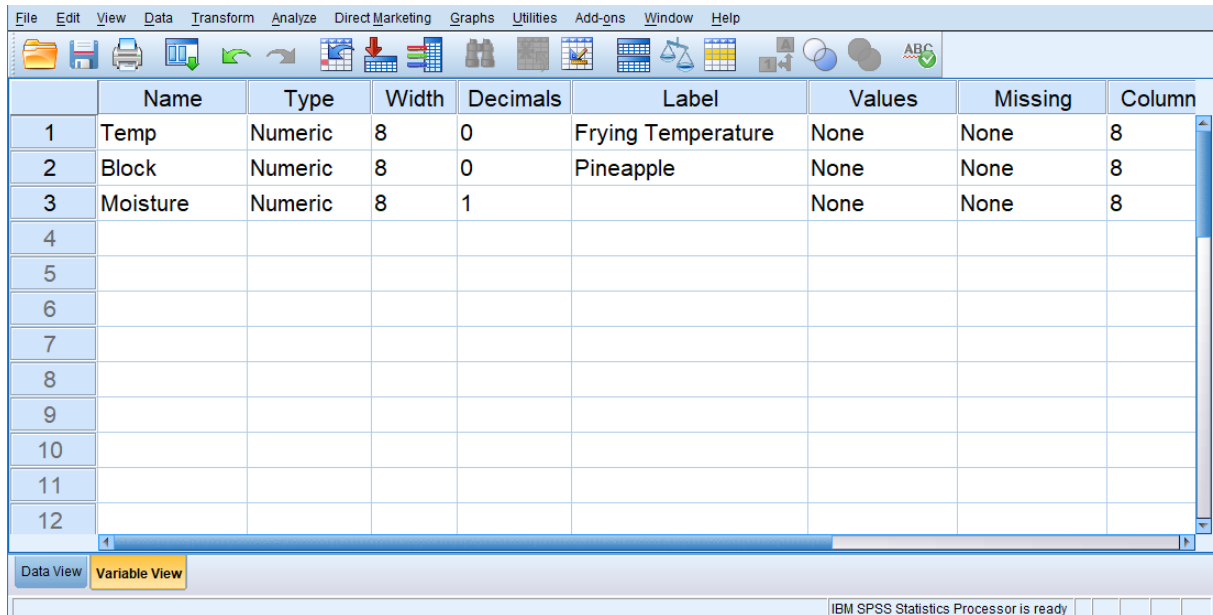
H_1 : ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดอย่างน้อย 2 บล็อกแตกต่างกัน

9.6.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลปริมาณความชื้นสับปะรดทอดในแต่ละอุณหภูมิที่ใช้ทอด จะต้องใช้ตัวแปร var จำนวน 3 ตัวแปรสำหรับป้อนข้อมูล 3 ตัว ได้แก่ อุณหภูมิในการทอด, บล็อก และ ปริมาณความชื้น ดังนั้นกำหนดชื่อตัวแปรเป็น **Temp**, **Block** และ **Moisture** ในหน้าต่าง Variable View ดังภาพที่ 9.2 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

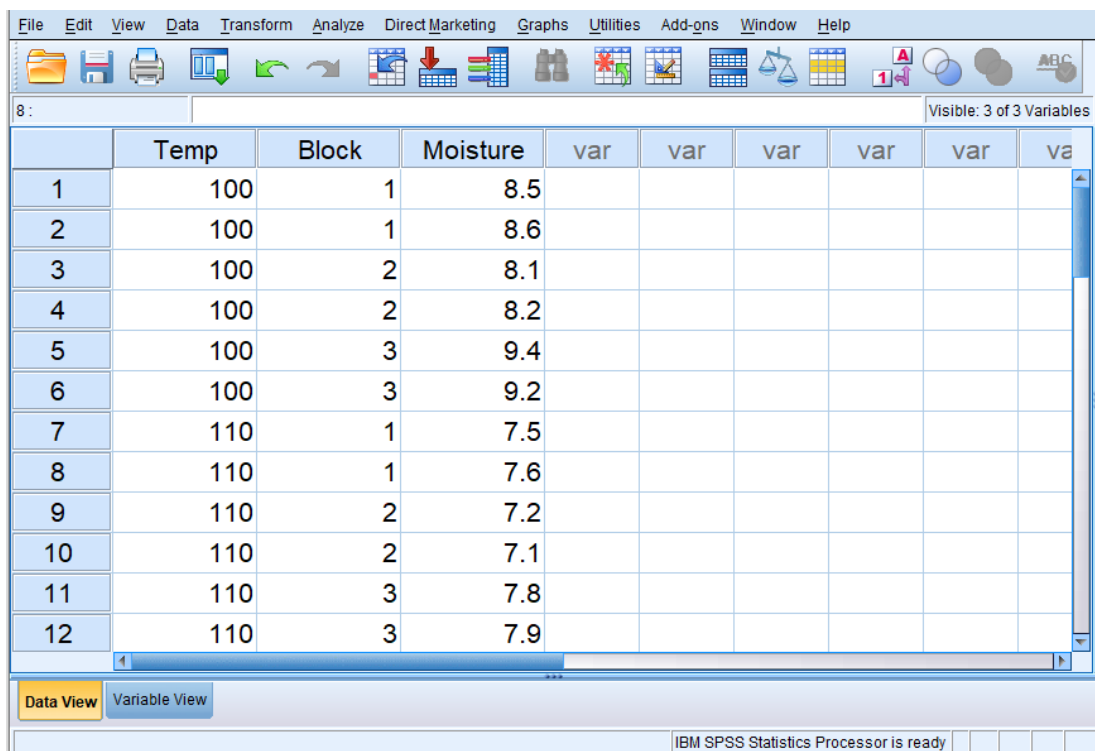
ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
อุณหภูมิในการทอด	Temp	Frying Temperature	0	ไม่ต้องกำหนด
บล็อก	Block	Pineapple	0	ไม่ต้องกำหนด
ปริมาณความชื้น	Moisture	ไม่ต้องกำหนด	1	ไม่ต้องกำหนด

หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



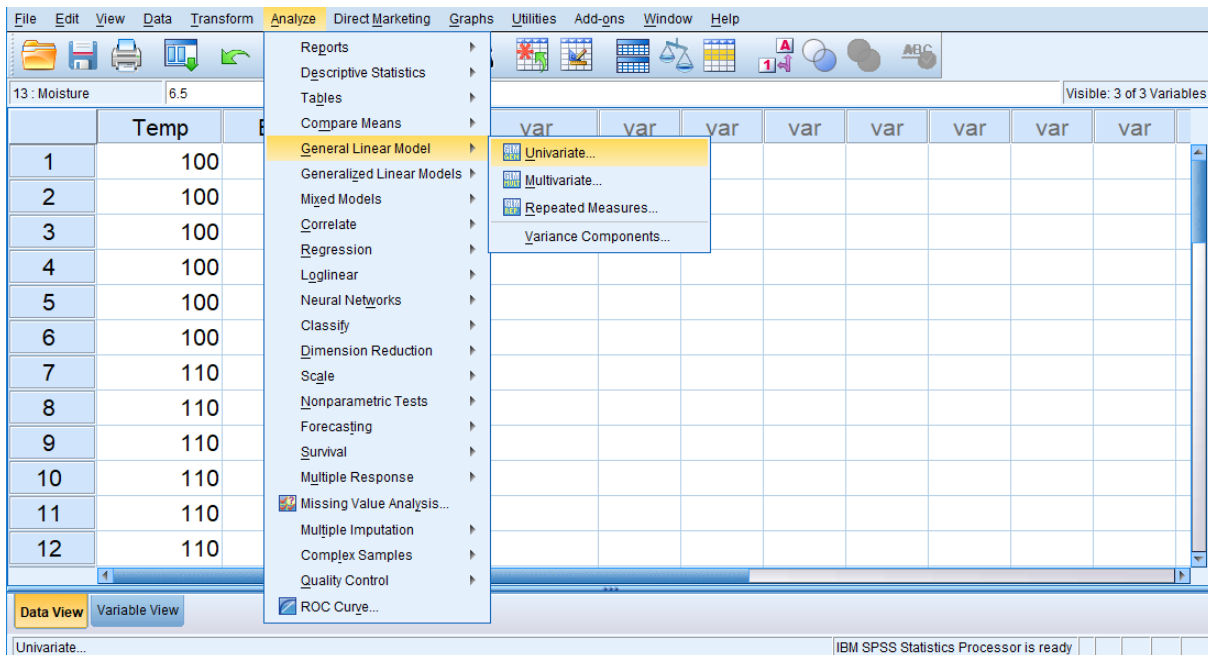
ภาพที่ 9.2 กำหนดรายละเอียดตัวแปร Temp, Block และ Moisture ในหน้าต่าง Variable View

(2) ป้อนข้อมูลปริมาณความชื้นสับปะรดทอดเรียงตามอุณหภูมิที่ใช้ทอดและตาม Block ในหน้าต่าง Data View ดังภาพที่ 9.3

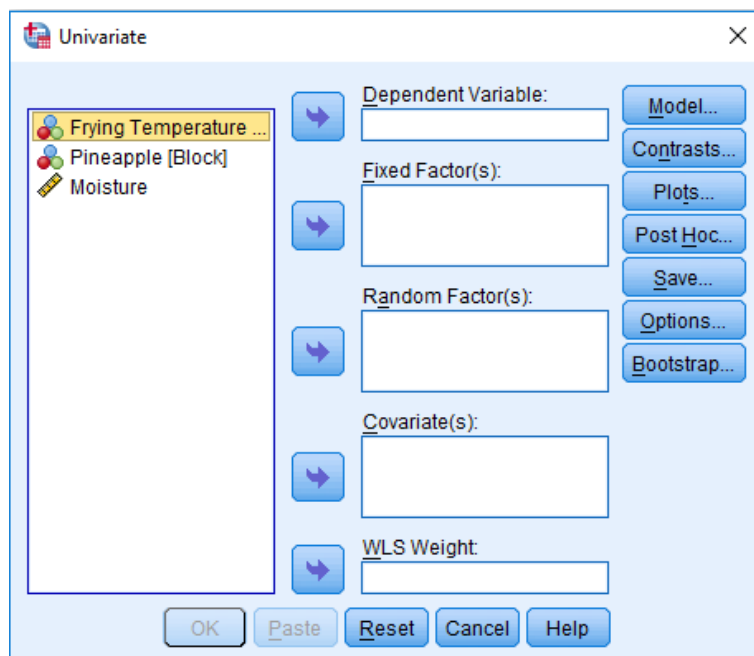


ภาพที่ 9.3 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 9.4 จะปรากฏหน้าจอ ดังภาพที่ 9.5

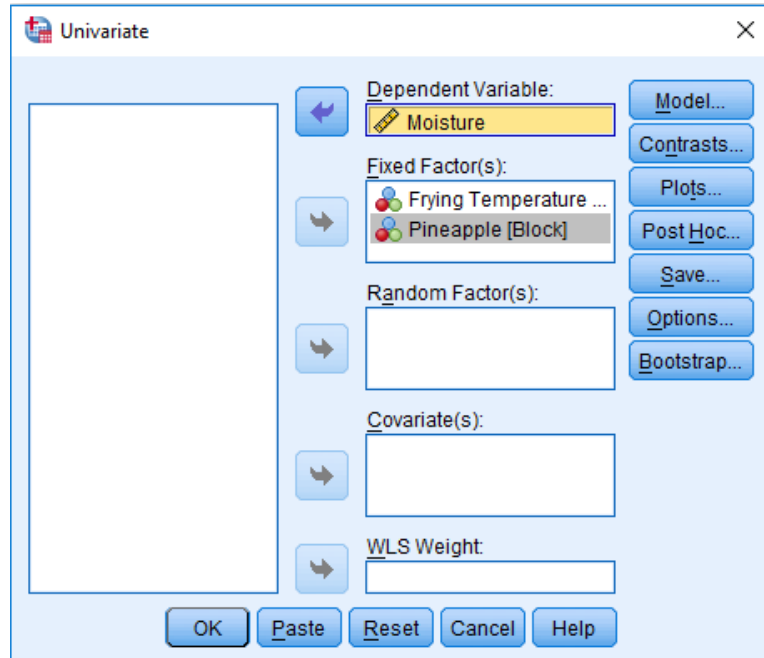


ภาพที่ 9.4 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



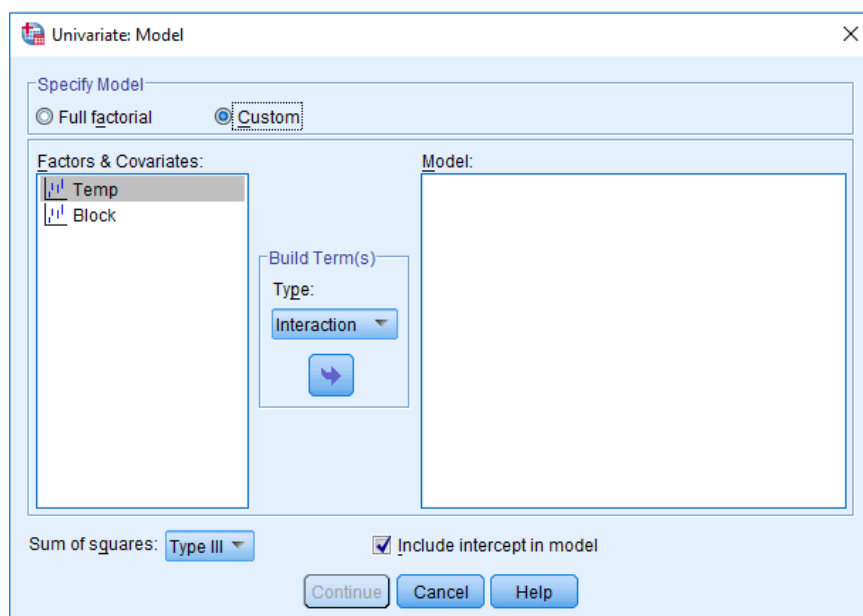
ภาพที่ 9.5 หน้าจอคำสั่ง Univariate

(4) เลือก ตัวแปร Moisture (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร Frying Temperature [Temp] กับ ตัวแปร Pineapple [Block] (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 9.6 ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate

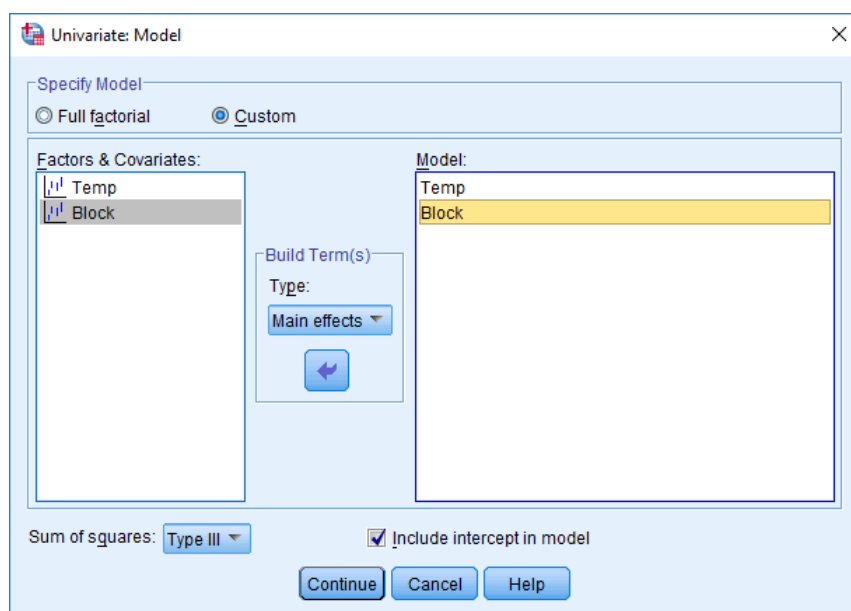


ภาพที่ 9.6 เลือกตัวแปร Moisture ใส่ในช่อง Dependent Variable และ เลือกตัวแปร Frying Temperature [Temp] กับ ตัวแปร Pineapple [Block] ใส่ในช่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Model...** ในการวางแผนการทดลองแบบ RCBD ต้องกำหนดตัวแบบ (Model) ชนิด **Custom** (ดังภาพที่ 9.7) จากนั้นให้กดเลือก **Type** ของ **Build Terms(s)** เป็น **Main effects** แล้วกดเลือก ตัวแปร **Temp** และ ตัวแปร **Block** จากกล่องซ้ายมือ (Factors & Covariates) ไปใส่ในกล่องขวามือ (Model) โดยคลิกเลือกทีละตัวแปร (ดังภาพที่ 9.8) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate (ข้อสังเกต จะไม่กำหนด Build Term(s) แบบ Interaction เนื่องจากเงื่อนไขของแผนการทดลองแบบ RCBD ต้องไม่มีปฏิสัมพันธ์ระหว่างสิ่งทดลองและบล็อก)

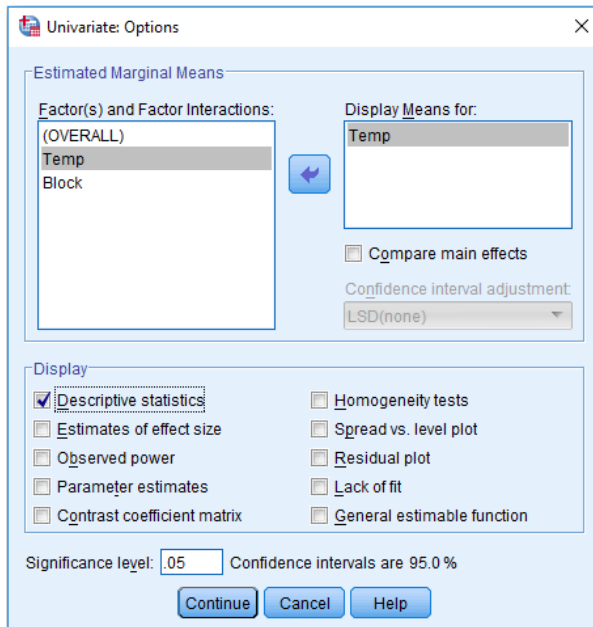


ภาพที่ 9.7 หน้าจอคำสั่ง Univariate : Model



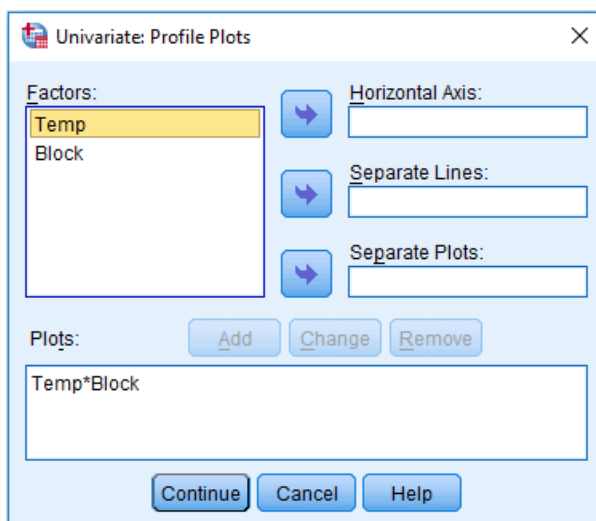
ภาพที่ 9.8 การเลือกตัวแปรในกล่อง Model

(6) เลือกคำสั่ง **Options...** จะได้หน้าจอตั้งภาพที่ 9.9 ให้เลือก **ตัวแปร Temp** จากกล่องซ้ายมือ (Factor(s) and Factor Interactions) ไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



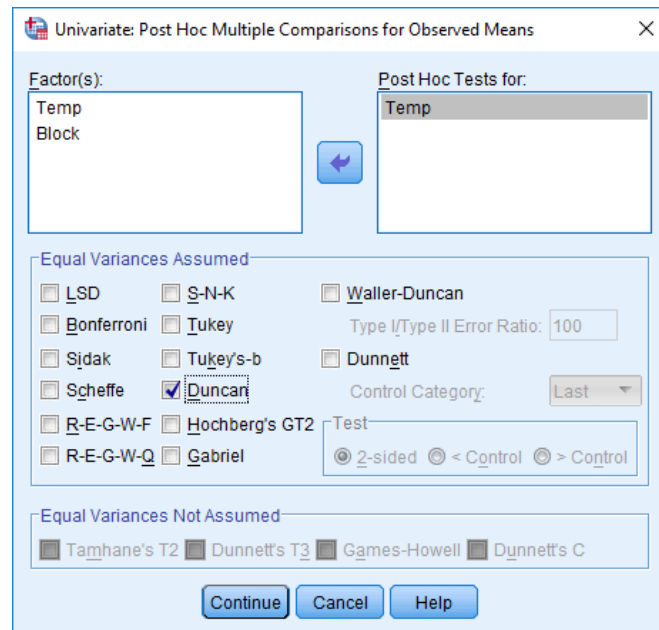
ภาพที่ 9.9 หน้าจอคำสั่ง Univariate : Options

(7) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลของสิ่งทดลองและบล็อก จะได้หน้าจอตั้งภาพที่ 9.10 กดเลือกตัวแปร Temp ในกล่อง Horizontal Axis (แกนนอน) และตัวแปร Block ในกล่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร Temp* Block ในกล่องด้านล่างเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 9.10 หน้าจอคำสั่ง Univariate : Profile Plots

(8) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 9.11 ให้กดเลือก **ตัวแปร Temp** จากกล่องซ้ายมือ (Factor(s)) ไปใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธีอื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 9.12



ภาพที่ 9.11 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

Univariate Analysis of Variance

Between-Subjects Factors

		N
Frying Temperature	100	6
	110	6
	120	6
Pineapple	1	6
	2	6
	3	6

Descriptive Statistics

Dependent Variable: Moisture

Frying Temperature	Pineapple	Mean	Std. Deviation	N
100	1	8.550	.0707	2
	2	8.150	.0707	2
	3	9.300	.1414	2
	Total	8.667	.5279	6
110	1	7.550	.0707	2
	2	7.150	.0707	2
	3	7.850	.0707	2
	Total	7.517	.3189	6
120	1	5.150	.0707	2
	2	5.750	.0707	2
	3	5.500	.1414	2
	Total	5.467	.2805	6
Total	1	7.083	1.5639	6
	2	7.017	1.0797	6
	3	7.550	1.7178	6
	Total	7.217	1.4106	18

Tests of Between-Subjects Effects

Dependent Variable: Moisture

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	32.543 ^a	4	8.136	82.522	.000
Intercept	937.445	1	937.445	9508.545	.000
Temp	31.530	2	15.765	159.905	.000
Block	1.013	2	.507	5.139	.023
Error	1.282	13	.099		
Total	971.270	18			
Corrected Total	33.825	17			

a. R Squared = .962 (Adjusted R Squared = .950)

Estimated Marginal Means

Frying Temperature

Dependent Variable: Moisture

Frying Temperature	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
100	8.667	.128	8.390	8.944
110	7.517	.128	7.240	7.794
120	5.467	.128	5.190	5.744

Post Hoc Tests

Frying Temperature

Homogeneous Subsets

Moisture

Duncan^{a, b}

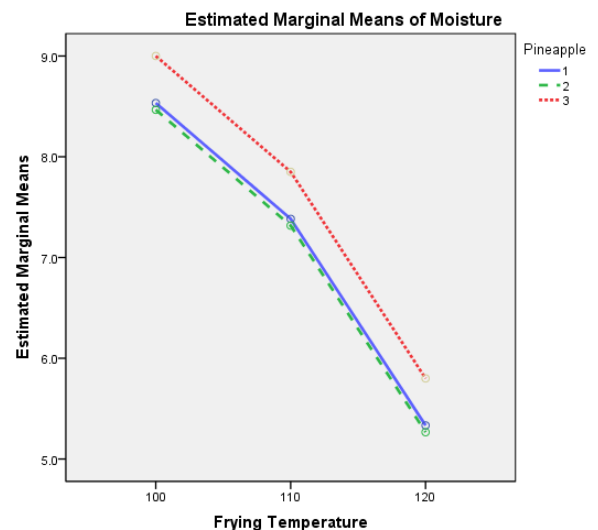
Frying Temperature	N	Subset		
		1	2	3
120	6	5.467		
110	6		7.517	
100	6			8.667
Sig.		1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .099.

a. Uses Harmonic Mean Sample Size = 6.000.

b. Alpha = .05.

Profile Plots



ภาพที่ 9.12 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนด้วยคำสั่ง Univariate

9.6.4 การอ่านผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ RCBD

จากผลการวิเคราะห์ในภาพที่ 9.12 แสดงผลลัพธ์ได้ 6 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Between-Subjects Factors เป็นส่วนที่แสดงรายละเอียดของข้อมูลที่ป้อน ดังนี้

Frying Temperature	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 3 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการทอด ได้แก่ 100, 110 และ 120 °C
Pineapple	คือ	บล็อกหรือสับปะรด ในที่นี้มี 3 บล็อก คือ สับปะรดลูกที่ 1, 2 และ 3
N	คือ	จำนวนข้อมูลในแต่ละสิ่งทดลองและแต่ละบล็อก ในที่นี้แต่ละสิ่งทดลองและแต่ละบล็อกมีจำนวนข้อมูลเท่ากัน คือ 6

ส่วนที่ 2 ตาราง Descriptive Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบจากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Frying Temperature	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 3 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการทอด ได้แก่ 100, 110 และ 120 °C
Pineapple	คือ	บล็อกหรือสับปะรด ในที่นี้มี 3 บล็อก คือ สับปะรดลูกที่ 1, 2 และ 3
Mean	คือ	ค่าเฉลี่ยความชื้นของแต่ละสิ่งทดลอง
Std. Deviation	คือ	ค่าส่วนเบี่ยงเบนมาตรฐานของค่าความชื้น แสดงการกระจายของข้อมูล
N	คือ	จำนวนข้อมูลในแต่ละสิ่งทดลองและแต่ละบล็อก

ส่วนที่ 3 ตาราง Tests of Between-Subjects Effects เป็นส่วนที่แสดงค่าสถิติต่างๆ ของตารางวิเคราะห์ความแปรปรวน เพื่อใช้ในการทดสอบสมมติฐานทางสถิติด้วยค่าสถิติ F หรือค่าความน่าจะเป็น Sig. เพื่อใช้ทดสอบสมมติฐาน อธิบายได้ ดังนี้

Type III Sum of Squares, df, Mean Square, F	คือ	ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในที่นี้จะพิจารณาค่า Sig. ของตัวแปร Temp และตัวแปร Block ตามสมมติฐานที่เขียนไว้ ดังนี้ (1) พิจารณาอิทธิพลของสิ่งทดลอง พบว่า ค่า Sig. ของตัวแปร Temp มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดอย่างน้อย 2 สิ่งทดลองแตกต่างกันที่ระดับนัยสำคัญ 0.05

- (2) พิจารณาอิทธิพลของบล็อก พบว่าค่า Sig. ของตัวแปร Block มีค่าเท่ากับ 0.023 ซึ่งน้อยกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดอย่างน้อย 2 บล็อกแตกต่างกันที่ระดับนัยสำคัญ 0.05

ส่วนที่ 4 ตาราง Estimated Marginal Means เป็นส่วนที่แสดงค่าเฉลี่ยโดยประมาณของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Frying Temperature	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 3 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการทอด ได้แก่ 100, 110 และ 120 °C
Mean	คือ	ค่าเฉลี่ยปริมาณความชื้นของแต่ละสิ่งทดลองโดยประมาณ
Std. Error	คือ	ค่าความคลาดเคลื่อนมาตรฐานในแต่ละกลุ่ม
95% Confidence Interval	คือ	ค่าที่แสดงขอบเขตบนล่างของค่าเฉลี่ยในแต่ละกลุ่มในช่วงความเชื่อมั่น 95%

ส่วนที่ 5 ตาราง Post Hoc Tests เป็นส่วนที่แสดงค่าสถิติสำหรับทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Tests) ซึ่งส่วนนี้จะนำมาใช้พิจารณาต่อเมื่อในส่วนที่ 3 ตาราง Tests of Between-Subjects Effects มีการปฏิเสธสมมติฐานหลัก (H_0) ของตัวแปร Temp อธิบายได้ ดังนี้

Frying Temperature	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกสิ่งทดลอง ในที่นี้มี 3 สิ่งทดลองจำแนกตามระดับอุณหภูมิในการทอด ได้แก่ 100, 110 และ 120 °C
N	คือ	จำนวนข้อมูลในแต่ละสิ่งทดลอง ในที่นี้แต่ละสิ่งทดลองมีจำนวนข้อมูลเท่ากับ 6
Subset	คือ	แสดงการแบ่งกลุ่มของสิ่งทดลอง ให้พิจารณารายละเอียดของ Subset ดังนี้ (1) มีจำนวน 3 Subset แต่ละ Subset มีสมาชิก 1 สิ่งทดลอง นั้นแสดงว่าแต่ละสิ่งทดลองมีค่าเฉลี่ยปริมาณความชื้นแตกต่างกันอย่างมีนัยสำคัญที่ระดับ 0.05 (2) ตัวเลขที่แสดงในแต่ละ Subset คือ ค่าเฉลี่ยปริมาณความชื้นของระดับอุณหภูมินั้นๆ เช่น Subset ที่ 1 ตัวเลข 5.467 แสดงค่าเฉลี่ยปริมาณความชื้นของสับปะรดที่ทอดอุณหภูมิ 120 °C (3) แต่ละ Subset จะใช้ตัวอักษรภาษาอังกฤษพิมพ์เล็กในการบ่งชี้กลุ่ม โดยจะนำตัวอักษรนี้ไปแสดงหลัง ค่าเฉลี่ย $\pm$ ค่าส่วนเบี่ยงเบนมาตรฐาน เช่น 5.47 $\pm$ 0.28c กรณีนี้มี 3 Subset จึงใช้ตัวอักษร a, b และ c บ่งชี้สมาชิกที่อยู่ใน Subset

ที่ 3, 2 และ 1 ตามลำดับ (เรียงลำดับตัวอักษรตามค่าเฉลี่ยจากมากไปน้อย ทั้งนี้ นักวิจัยอาจเรียงจากน้อยไปมากได้เช่นกัน)

ส่วนที่ 6 Profile Plots เป็นส่วนที่แสดงกราฟเส้นของตัวแปรที่ต้องการทดสอบ

ในที่นี้คือตัวแปร Moisture กราฟที่แสดงใช้ค่าเฉลี่ยโดยประมาณ (จากตาราง Estimated Marginal Means) มาแสดงโดยจำแนกตามตัวแปรที่ได้เลือกในขั้นตอนการวิเคราะห์ คำสั่ง Plots... ในที่นี้ได้เลือกตัวแปร Temp เป็นตัวแปรแกนนอน จึงเห็นกราฟในแกนนอนมี 3 สเกล (100, 110 และ 120) และได้เลือกตัวแปร Block แสดงเป็นเส้นกราฟ จึงเห็นในกราฟมี 3 เส้น (แต่ละเส้นแทนสับปะรดแต่ละบล็อก) นักวิจัยสามารถใช้กราฟ Profile Plots ในการพิจารณาแนวโน้มค่าเฉลี่ยระหว่างระดับต่างๆ ของปัจจัยที่ศึกษา ดังตัวอย่างนี้ เมื่อเพิ่มอุณหภูมิในการทอดมีผลทำให้ปริมาณความชื้นของสับปะรดมีค่าลดลงไม่ว่าจะใช้สับปะรดจากลูกใด

9.6.5 การนำเสนอผลลัพธ์จากการวิเคราะห์ข้อมูลที่ได้จากการวางแผนการทดลองแบบ RCBD

จากภาพที่ 9.12 ในตาราง Tests of Between-Subjects Effects เมื่อใช้สถิติ F ทดสอบอิทธิพลของสิ่งทดลอง (Temp) พบว่า ค่า Sig. เท่ากับ 0.000 ซึ่งน้อยกว่า 0.05 นั่นคือ ปฏิเสธสมมติฐานหลัก (H_0) จึงสรุปได้ว่า ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดอย่างน้อย 2 สิ่งทดลองแตกต่างกัน ($p \leq 0.05$) และ เมื่อพิจารณาอิทธิพลของบล็อก (Block) พบว่า ค่า Sig. เท่ากับ 0.023 ซึ่งน้อยกว่า 0.05 นั่นคือ ปฏิเสธ H_0 จึงสรุปได้ว่า ค่าเฉลี่ยปริมาณความชื้นของสับปะรดทอดอย่างน้อย 2 บล็อกแตกต่างกัน ($p \leq 0.05$)

ในการเขียนผลการวิเคราะห์จะนำข้อมูลค่าเฉลี่ย (Mean) และค่าส่วนเบี่ยงเบนมาตรฐาน (Std. Deviation) ของแต่ละสิ่งทดลองจากข้อมูลในตาราง Descriptive Statistics และผลการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณจากข้อมูลในตาราง Post Hoc Tests มาใช้ในการเขียนผล ดังแสดงในตารางที่ 9.5

ตารางที่ 9.5 ค่าเฉลี่ยปริมาณความชื้นสับปะรดที่ทอดด้วยอุณหภูมิที่แตกต่างกัน

อุณหภูมิในการทอด (°C)	ปริมาณความชื้น (%)
100	$8.67 \pm 0.53a$
110	$7.52 \pm 0.32b$
120	$5.47 \pm 0.28c$

^{a-c}ตัวอักษรที่แตกต่างกันในแนวดิ่งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

9.7 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD

9.7.1 ตัวอย่างข้อมูล

ตามที่สำนักงานคณะกรรมการอาหารและยา (อย.) มีนโยบายจะยกเว้นประกาศ ห้ามใช้ไขมันทรานส์ในการผลิตอาหารเพื่อลดความเสี่ยงโรคหัวใจและหลอดเลือด ทำให้ผู้ประกอบการต้องเร่งทำการปรับปรุงสูตรผลิตภัณฑ์ของตัวเองโดยเฉพาะผลิตภัณฑ์ขนมอบที่มีการใช้ไขมันทรานส์เพื่อรองรับประกาศดังกล่าว นักพัฒนาผลิตภัณฑ์ของบริษัทผู้ผลิตเค้กแห่งหนึ่งจึงต้องการหาชนิดไขมันหรือน้ำมันจากพืชชนิดที่ให้กรดไขมันไม่อิ่มตัวมาใช้แทนที่เนยในผลิตภัณฑ์เค้ก โดยมีเป้าหมายคือผู้บริโภคให้คะแนนความชอบไม่แตกต่างกับสูตรที่ใช้เนยในการผลิต

นักพัฒนาผลิตภัณฑ์ได้วางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) โดยใช้แหล่งของไขมันที่แตกต่างกัน 3 ชนิดในปริมาณที่เท่ากันผลิตภัณฑ์ ได้แก่ เนย (สูตรควบคุม), น้ำมันรำข้าว และ น้ำมันถั่วเหลือง ให้ผู้ทดสอบชิมจำนวน 30 คน ให้คะแนนความชอบเค้กทั้งสามสิ่งทดลองโดยใช้สเกล 9-point hedonic scale (1=ไม่ชอบมากที่สุด และ 9= ชอบมากที่สุด) ทั้งนี้ในการสุ่มตัวอย่างเค้กได้ทำการสุ่มลำดับการสุ่มให้แก่วัดทดสอบแต่ละคนเพื่อลดผลกระทบที่อาจเกิดขึ้นจากลำดับการชิม ได้ผลการประเมินความชอบดังนี้

ผู้ทดสอบ คนที่	ชนิดไขมัน		
	เนย	น้ำมัน รำข้าว	น้ำมัน ถั่วเหลือง
1	8	6	6
2	8	6	7
3	8	6	7
4	8	6	6
5	7	6	6
6	7	6	6
7	7	7	7
8	7	7	7
9	8	7	6
10	8	8	6
11	8	7	7
12	8	7	7
13	6	7	7
14	6	7	5
15	6	7	5

ผู้ทดสอบ คนที่	ชนิดไขมัน		
	เนย	น้ำมัน รำข้าว	น้ำมัน ถั่วเหลือง
16	8	7	6
17	8	7	6
18	7	7	6
19	7	7	6
20	7	8	7
21	7	8	7
22	7	8	7
23	7	8	6
24	7	8	6
25	7	8	7
26	8	6	7
27	6	6	6
28	6	6	6
29	6	8	7
30	6	8	7

จากข้อมูลดังกล่าว สามารถสรุปรายละเอียดของข้อมูล ได้ดังนี้

- (1) ปัจจัยที่ศึกษา มี 1 ปัจจัย คือ ชนิดของไขมัน
- (2) จำนวนสิ่งทดลอง (Treatment) เท่ากับ 3 สิ่งทดลอง (เท่ากับจำนวนชนิดไขมัน) โดยเค้กที่ใช้เนย คือ สิ่งทดลองควบคุม
- (3) บล็อก (Block) คือ ผู้ทดสอบชิม และ จำนวนบล็อก เท่ากับ 30 (เท่ากับจำนวนผู้ทดสอบชิม)
- (4) จำนวนหน่วยทดลอง (Experimental unit) ในแต่ละสิ่งทดลอง เท่ากับ 30 หน่วย
- (5) จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 90 หน่วย
- (6) ค่าสังเกต (ตัวแปรตาม) คือ ค่าคะแนนความชอบ

9.7.2 การเขียนสมมติฐานสำหรับการวิเคราะห์ความแปรปรวนจากแผนการทดลองแบบ RCBD

นักวิจัยสามารถเขียนสมมติฐานสำหรับทดสอบความแตกต่างของค่าเฉลี่ยได้ 2 ข้อ ดังนี้

สมมติฐานข้อ 1 ทดสอบอิทธิพลของสิ่งทดลองที่มีผลต่อค่าคะแนนความชอบ

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบไม่ขึ้นกับชนิดของไขมัน

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบขึ้นกับชนิดของไขมัน

สมมติฐานข้อ 2 ทดสอบอิทธิพลของบล็อกที่มีผลต่อค่าคะแนนความชอบ

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบแต่ละคนไม่แตกต่างกัน

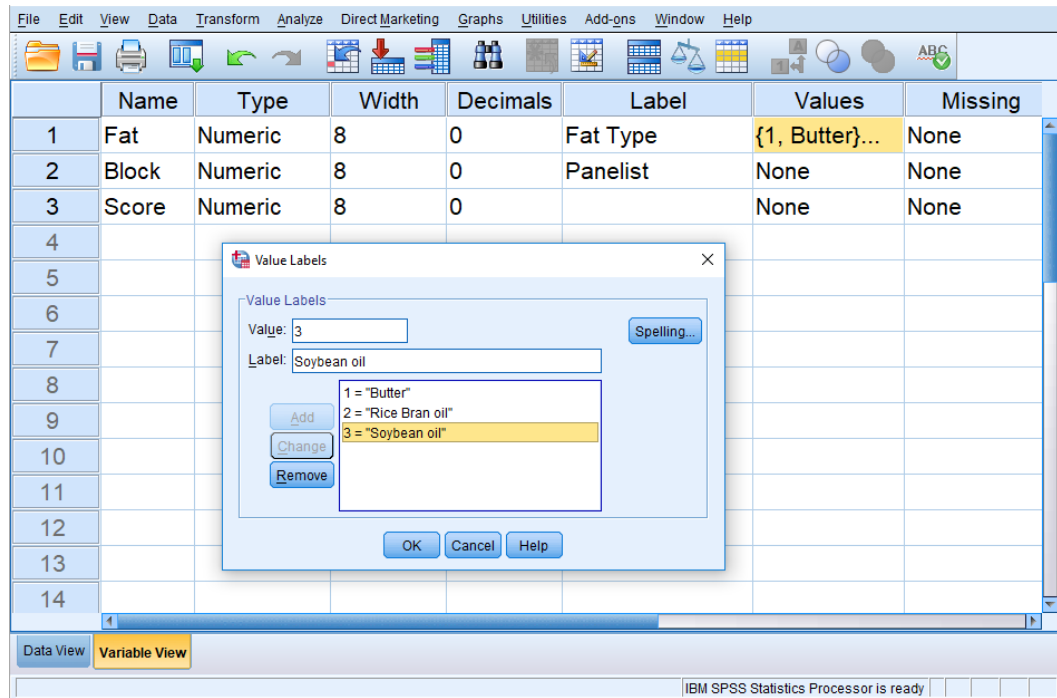
H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบอย่างน้อย 2 คนแตกต่างกัน

9.7.3 การใช้คำสั่งในโปรแกรม SPSS ในการวิเคราะห์ความแปรปรวนแบบสองทาง

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลคะแนนความชอบเค้กแต่ละสิ่งทดลอง ให้กำหนดรายละเอียดตัวแปรในหน้าต่าง Variable View ดังภาพที่ 9.13 ในกรณีนี้ใช้ตัวแปร var จำนวน 3 ตัวแปร สำหรับป้อนข้อมูล 3 ตัว ได้แก่ ชนิดไขมัน, ผู้ทดสอบ และ คะแนนความชอบ ดังนั้นกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ชนิดไขมัน	Fat	Fat Type	0	1 = Butter 2 = Rice Bran oil 3 = Soybean oil
ผู้ทดสอบ	Block	Panelist	0	ไม่ต้องกำหนด
คะแนนความชอบ	Score	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด

หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



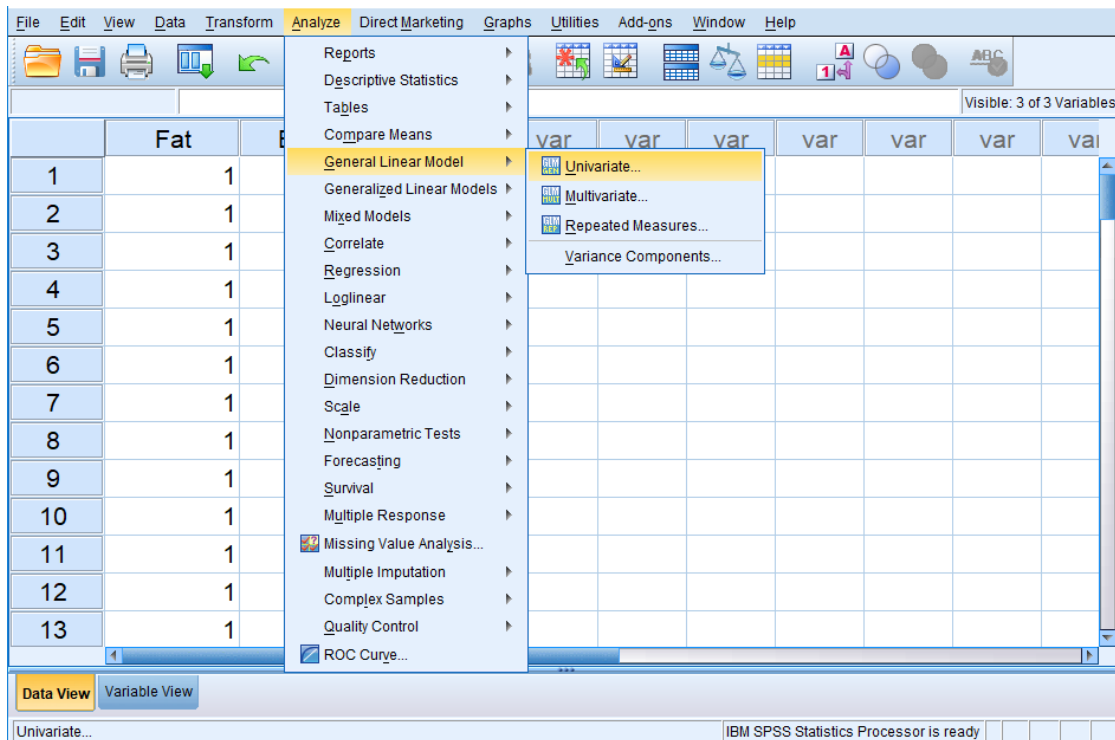
ภาพที่ 9.13 กำหนดรายละเอียดตัวแปรในหน้าต่าง Variable View

(2) ป้อนข้อมูลชนิดไขมัน ลำดับของผู้ทดสอบ (บล็อก) และคะแนนความชอบ ในหน้าต่าง Data View ดังภาพที่ 9.14 โดยป้อนข้อมูลเรียงตามชนิดไขมันที่ใช้และลำดับของผู้ทดสอบ

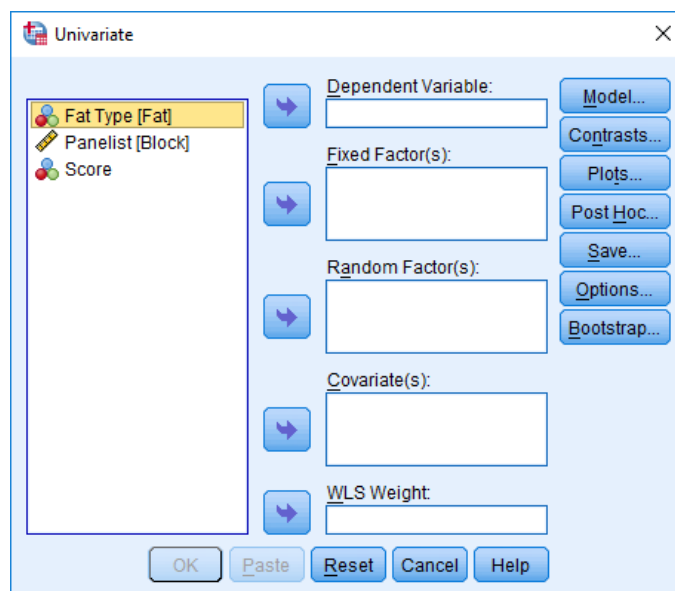
15:										
	Fat	Block	Score	var	var	var	var	var	var	var
1	1	1	8							
2	1	2	8							
3	1	3	8							
4	1	4	8							
5	1	5	7							
6	1	6	7							
7	1	7	7							
8	1	8	7							
9	1	9	8							
10	1	10	8							
11	1	11	8							
12	1	12	8							
13	1	13	6							

ภาพที่ 9.14 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze ► General Linear Model ► Univariate...** ดังภาพที่ 9.15 จะปรากฏหน้าจอ ดังภาพที่ 9.16 จะเห็นว่าในกล่องซ้ายมือจะแสดงชื่อของตัวแปรแต่ละตัว

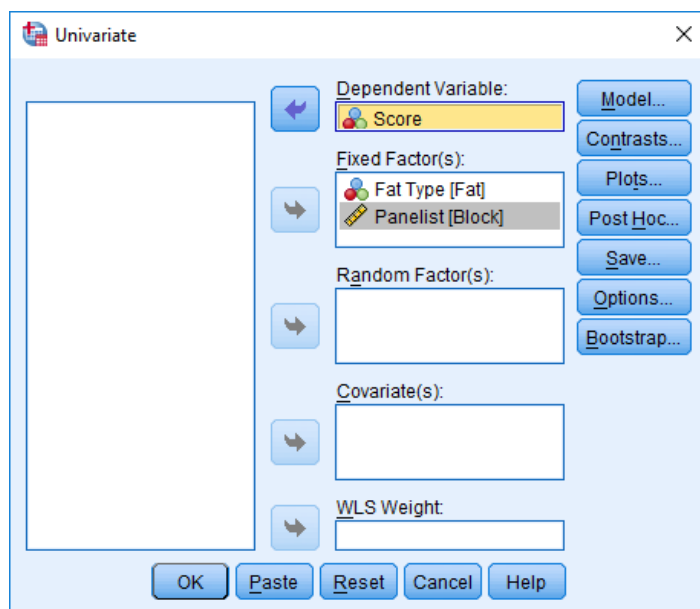


ภาพที่ 9.15 เมนูหน้าจอคำสั่ง Analyze ► General Linear Model ► Univariate...



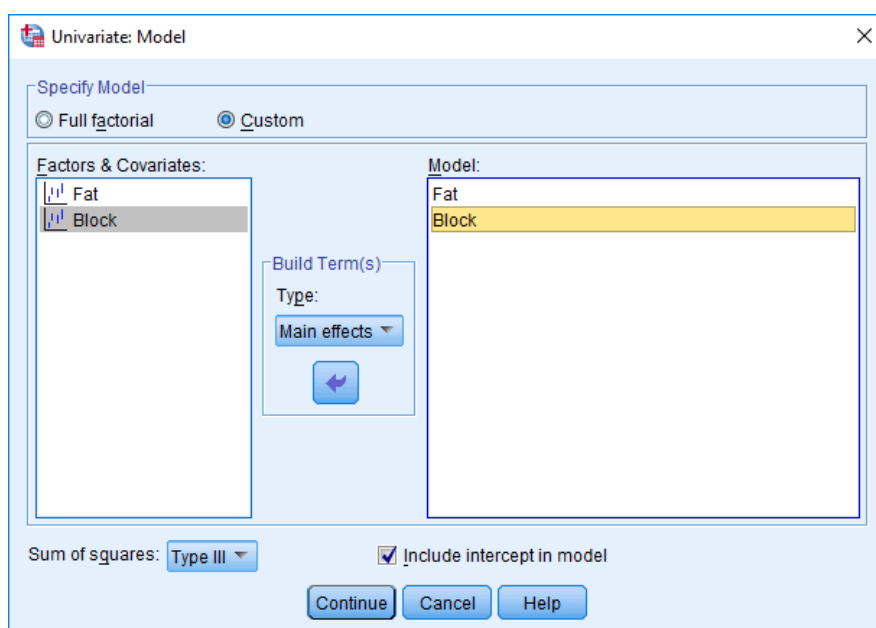
ภาพที่ 9.16 หน้าจอคำสั่ง Univariate

(4) เลือก ตัวแปร **Fat Type [Fat]** และ ตัวแปร **Panelist [Block]** (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** และเลือก ตัวแปร **Score** (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** จะได้ดังภาพที่ 9.17 ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



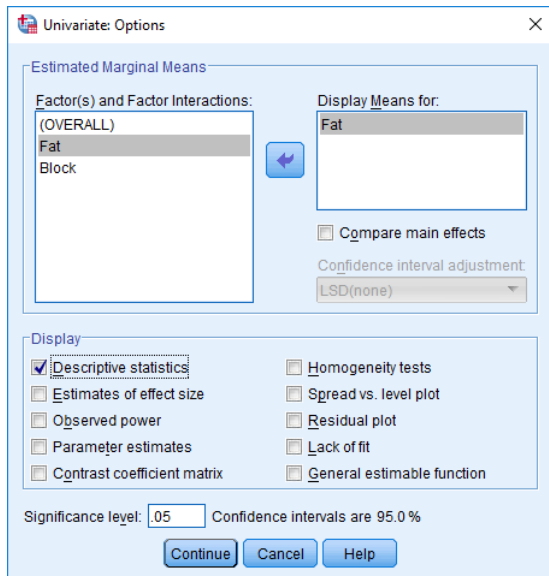
ภาพที่ 9.17 เลือกตัวแปร Score ใส่ในกล่อง Dependent Variable และตัวแปร Fat Type [Fat] และตัวแปร Panelist [Block] ใส่ในกล่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Model...** กำหนดตัวแบบ (Model) ชนิด **Custom** จากนั้นให้กดเลือก **Type** ของ **Build Terms(s)** เป็น **Main effects** แล้วกดเลือก ตัวแปร **Fat** และ ตัวแปร **Block** จากกล่องซ้ายมือ (Factors & Covariates) ไปใส่ในกล่องขวามือ (Model) โดยคลิกเลือกทีละตัวแปร (ดังภาพที่ 9.18) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



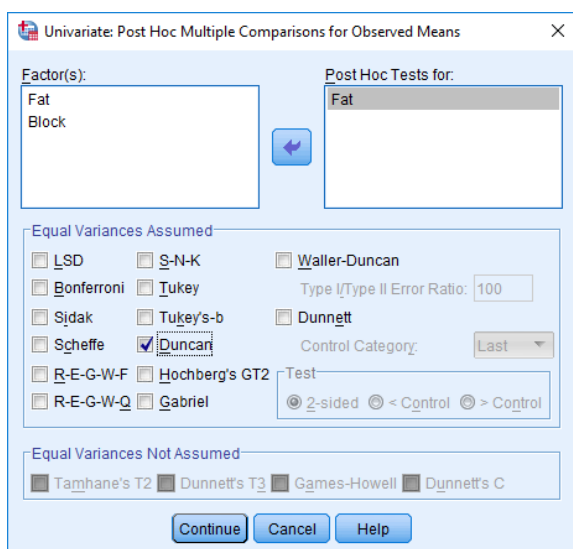
ภาพที่ 9.18 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Options...** จะได้หน้าจอตั้งภาพที่ 9.19 ให้กดเลือก **ตัวแปร Fat** จากกล่องซ้ายมือ ไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและ ส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของ คำสั่ง Univariate



ภาพที่ 9.19 หน้าจอคำสั่ง Univariate : Options

(7) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการทดสอบแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ในกรณีที่มีค่าที่มีความแตกต่างของค่าเฉลี่ยอย่างน้อย 1 คู่ จะได้หน้าจอตั้งภาพที่ 9.20 ให้กดเลือก **ตัวแปร Fat** ในกล่องซ้ายมือ (Factor(s)) ใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธีอื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 9.21



ภาพที่ 9.20 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

Univariate Analysis of Variance

Between-Subjects Factors

	Value Label	N
Fat Type	1 Butter	30
	2 Rice Bran oil	30
	3 Soybean oil	30
Panelist	1	3
	2	3
	3	3
	4	3
	5	3
	6	3
	7	3

Rows 1 through 10 of 33



Descriptive Statistics

Dependent Variable: Score

Fat Type	Panelist	Mean	Std. Deviation	N
Butter	1	8.00	.	1
	2	8.00	.	1
	3	8.00	.	1
	4	8.00	.	1
	5	7.00	.	1
	6	7.00	.	1
	7	7.00	.	1
	8	7.00	.	1
	9	8.00	.	1
	10	8.00	.	1

Rows 1 through 10 of 124 (some rows are hidden)



Tests of Between-Subjects Effects

Dependent Variable: Score

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	25.644 ^a	31	.827	1.590	.064
Intercept	4216.178	1	4216.178	8103.258	.000
Fat	9.156	2	4.578	8.798	.000
Block	16.489	29	.569	1.093	.378
Error	30.178	58	.520		
Total	4272.000	90			
Corrected Total	55.822	89			

a. R Squared = .459 (Adjusted R Squared = .170)

Estimated Marginal Means

Fat Type

Dependent Variable: Score

Fat Type	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Butter	7.133	.132	6.870	7.397
Rice Bran oil	7.000	.132	6.736	7.264
Soybean oil	6.400	.132	6.136	6.664

Post Hoc Tests

Fat Type

Homogeneous Subsets

Score

Duncan^{a,b}

Fat Type	N	Subset	
		1	2
Soybean oil	30	6.40	
Rice Bran oil	30		7.00
Butter	30		7.13
Sig.		1.000	.477

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .520.

a. Uses Harmonic Mean Sample Size = 30.000.
b. Alpha = .05.

ภาพที่ 9.21 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนด้วยคำสั่ง Univariate

จากภาพที่ 9.21 เมื่อพิจารณาในตาราง Between-Subjects Factors และตาราง Descriptive Statistics จะพบว่าในตารางแสดงข้อมูลไม่ครบ ให้นักวิจัยสังเกตที่ลูกศรใต้ตารางและคำอธิบายใต้ตารางที่บ่งชี้ว่ายังมีข้อมูลที่ซ่อนอยู่ไม่ได้แสดง เช่น ในตาราง Between-Subjects Factors แสดงข้อมูล 10 ตัวจากทั้งหมด 33 ตัว ดังนั้นหากนักวิจัยต้องการเรียกดูข้อมูลทั้งหมดทำได้โดยกดคลิกที่ตารางนั้นจะปรากฏลูกศรสีแดงข้างซ้ายของตารางนั้นดังภาพที่ 9.22 ขั้นตอนถัดไปให้กดเมาส์คลิกขวาจะปรากฏกล่องคำสั่งให้เลือกคำสั่ง **Set Rows to Display...** แล้วป้อนจำนวนข้อมูลที่ต้องการให้แสดง เช่น ในกรณีนี้ต้องการให้แสดงทั้งหมด 33 ตัว

Univariate Analysis of Variance

The image shows two SPSS output tables. The top table, 'Between-Subjects Factors', lists 'Fat Type' and 'Panelist' with their respective 'Value Label' and 'N'. A red arrow points to the first row of the 'Panelist' section. A right-click context menu is open over the table, with 'Set Rows to Display...' highlighted. The bottom table, 'Descriptive Statistics', shows the mean and standard deviation for each 'Fat Type' across 'Panelist' groups. A right-click context menu is also open over this table, with 'Set Rows to Display...' highlighted. A dialog box titled 'Set Rows to Display' is shown, with 'Rows to Display' set to 33 and 'Widow/Orphan Tolerance' set to 1.

	Value Label	N
Fat Type 1	Butter	30
2	Rice Bran oil	30
3	Soybean oil	30
Panelist 1		3
2		3
3		3
4		3
5		3
6		3
7		3

Rows 1 through 10 of 33

Fat Type	Panelist	Mean	Std. Deviation
Butter	1	8.00	.00
	2	8.00	.00
	3	8.00	.00
	4	8.00	.00
	5	7.00	.00
	6	7.00	.00
	7	7.00	.00
	8	7.00	.00
	9	8.00	.00
	10	8.00	.00

Rows 1 through 10 of 124 (some rows are hidden)

Set Rows to Display dialog box:

Rows to Display: 33
Widow/Orphan Tolerance: 1
OK Cancel Help

ภาพที่ 9.22 การใช้คำสั่ง Set Rows to Display... เพื่อแสดงข้อมูลที่ซ่อนอยู่

9.7.4 การสรุปผลลัพธ์การวิเคราะห์จากการวางแผนการทดลองแบบ RCBD

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่ได้กำหนดไว้ ดังนี้

สมมติฐานข้อ 1 ทดสอบอิทธิพลของสิ่งทดลองที่มีผลต่อค่าคะแนนความชอบ

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบไม่ขึ้นกับชนิดของไขมัน

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบขึ้นกับชนิดของไขมัน

จากภาพที่ 9.21 ในตาราง Tests of Between-Subjects Effects พบว่าค่า Sig. ของตัวแปร Fat มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบขึ้นกับชนิดของไขมัน ($p \leq 0.05$) เนื่องจากตัวแปร Fat มี 3 ระดับหรือ 3 สิ่งทดลอง เมื่อปฏิเสธสมมติฐานหลักจึงต้องพิจารณาผลการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณในตาราง Post Hoc Tests ซึ่งพบว่ามี 2 Subset โดยเค้กที่ใช้ไขมันรำข้าวมีค่าเฉลี่ยคะแนนความชอบไม่แตกต่างกับเค้กที่ใช้เนย (สิ่งทดลองควบคุม) เนื่องจากอยู่ใน Subset เดียวกัน

สมมติฐานข้อ 2 ทดสอบอิทธิพลของบล็อกที่มีผลต่อค่าคะแนนความชอบ

H_0 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบแต่ละคนไม่แตกต่างกัน

H_1 : ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบอย่างน้อย 2 คนแตกต่างกัน

จากภาพที่ 9.21 ในตาราง Tests of Between-Subjects Effects ค่า Sig. ของตัวแปร Block มีค่าเท่ากับ 0.378 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ค่าเฉลี่ยคะแนนความชอบเค้กของผู้ทดสอบแต่ละคนไม่แตกต่างกัน ($p > 0.05$)

ข้อสังเกต จากกรณีศึกษานี้เป็นกรณีประเมินคุณภาพทางประสาทสัมผัสโดยใช้การวางแผนการทดลองแบบ RCBD และกำหนดให้ผู้ทดสอบชิมเป็นบล็อก แม้ว่าไม่พบอิทธิพลของบล็อกที่มีต่อค่าคะแนนความชอบ ในการทดลองครั้งต่อไปนักวิจัยไม่สามารถจะสรุปได้ว่า ไม่จำเป็นต้องใช้การวางแผนการทดลองแบบ RCBD เพราะในบางกรณีแม้ว่าจะไม่มีอิทธิพลของบล็อกแต่มีความจำเป็นต้องลดความคลาดเคลื่อนดังกล่าวออกจาก สิ่งทดลอง ดังนั้นในการทดลองครั้งต่อไปจึงยังจำเป็นต้องใช้แผนการทดลองแบบ RCBD สำหรับการประเมินคุณภาพทางประสาทสัมผัส

9.7.5 การนำเสนอผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ RCBD

ในการเขียนผลการวิเคราะห์จะนำข้อมูลค่าเฉลี่ย (Mean) และค่าส่วนเบี่ยงเบนมาตรฐาน (Std. Deviation) ของแต่ละสิ่งทดลองจากข้อมูลในตาราง Descriptive Statistics และผลการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณจากข้อมูลในตาราง Post Hoc Tests มาใช้ในการเขียนผล ดังแสดงในตารางที่ 9.6

ตารางที่ 9.6 ค่าเฉลี่ยคะแนนความชอบเค้กที่ใช้ไขมันต่างชนิดกัน ประเมินโดยผู้ทดสอบจำนวน 30 คนด้วยสเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด)

ชนิดไขมัน	คะแนนความชอบ
เนย	7.1 ± 0.8a
น้ำมันรำข้าว	7.0 ± 0.8a
น้ำมันถั่วเหลือง	6.4 ± 0.6b

^{a-b}ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

เมื่อพิจารณาข้อมูลจากตารางที่ 9.6 นักพัฒนาผลิตภัณฑ์สามารถใช้น้ำมันรำข้าวแทนที่เนยในการผลิตเค้กเนื่องจากผู้ทดสอบให้คะแนนความชอบไม่แตกต่างกับสูตรที่ใช้เนยในการผลิตที่ระดับนัยสำคัญ 0.05



บทที่ 10

การทดลองแบบแฟกทอเรียล

ในบทที่ 8 และ 9 เป็นการพิจารณาการวางแผนการทดลองที่มีเพียง 1 ปัจจัย เช่น อุณหภูมิในการทอด ที่มีผลต่อการอมน้ำมันของมันฝรั่งทอด ทั้งนี้การศึกษาในบางเรื่องอาจจำเป็นต้องเกี่ยวข้องกับปัจจัยมากกว่า 1 ปัจจัย เช่น นักวิจัยคาดว่าอุณหภูมิและเวลาในการทอดต่างมีผลต่อการอมน้ำมันของมันฝรั่งทอด กรณีนี้ปัจจัยที่ต้องการศึกษาจึงมี 2 ปัจจัย คือ อุณหภูมิในการทอด และ เวลาในการทอด สำหรับการวางแผนการทดลองเพื่อทดสอบสมมติฐานงานวิจัยดังกล่าวสามารถทำได้ 2 กรณี คือ กรณีที่ 1 ทำการทดลองทีละปัจจัยโดยให้ปัจจัยใดปัจจัยหนึ่งคงที่ก่อน ซึ่งจะต้องทำการทดลองหลายครั้งและมีข้อเสียคือไม่สามารถศึกษาอิทธิพลร่วมระหว่างปัจจัยได้ (Interaction effects) และ กรณีที่ 2 ทำการทดลองแบบแฟกทอเรียล (Factorial design) ซึ่งเป็นการทดลองที่สามารถศึกษาอิทธิพลของปัจจัยหลักได้หลายๆ ปัจจัยพร้อมกัน และยังสามารถศึกษาอิทธิพลร่วมระหว่างปัจจัยได้

10.1 การทดลองแบบแฟกทอเรียล (Factorial experiment)

การทดลองแบบแฟกทอเรียลเป็นการทดลองที่สนใจศึกษาอิทธิพลของหลายๆ ปัจจัยพร้อมกัน โดยแต่ละปัจจัยมีหลายระดับ การจัดสิ่งทดลองหรือทรีทเมนต์ของการทดลองแบบนี้จะพิจารณาที่การจัดกลุ่ม (Treatment combination) ของระดับต่างๆ ของปัจจัยดังกล่าวที่เป็นไปได้ทั้งหมด ดังนั้นการทดลองแบบแฟกทอเรียลจึงมีประโยชน์สำหรับงานวิจัยที่ไม่ค่อยทราบเกี่ยวกับระดับที่เหมาะสมหรือระดับที่มีความสำคัญของปัจจัยต่างๆ การทดลองแบบนี้จะช่วยหาระดับที่เหมาะสมได้ นอกจากนั้นยังมีประโยชน์สำหรับใช้ในการศึกษาอิทธิพลของปัจจัยหนึ่งๆ ว่ามีผลต่อปัจจัยอื่นๆ ที่ระดับต่างๆ กัน หรือ ในสถานะที่ต่างกัน (Interaction) หรือไม่

10.2 การจัดกลุ่มสิ่งทดลอง (Treatment combination)

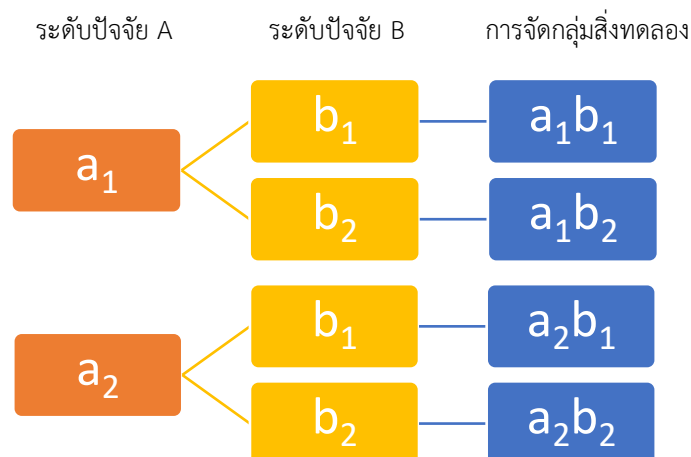
ในการทดลองแบบแฟกทอเรียลนั้นการจัดกลุ่มสิ่งทดลองทำได้โดยจัดกลุ่มระหว่างระดับต่างๆ ของปัจจัยที่ศึกษา โดยทั่วไปสัญลักษณ์ที่ใช้แทนระดับของปัจจัย คือ การใช้อักษรภาษาอังกฤษตัวพิมพ์เล็กและตัวเลขที่เขียนห้อยตัวอักษรเพื่อแสดงระดับของปัจจัยนั้น เช่น กรณีปัจจัยที่ศึกษาคือ ปัจจัย A และมี 2 ระดับ (ระดับต่ำ และ ระดับสูง) สามารถเขียนระดับของปัจจัยได้ ดังนี้ a_1 แทน ปัจจัย A ที่ระดับต่ำ และ a_2 แทน ปัจจัย A ที่ระดับสูง

การพิจารณาจัดกลุ่มสิ่งทดลองทำได้ 2 วิธี คือ การใช้ตาราง และ แผนภูมิแกว่งปลา โดยให้ยึดเกณฑ์ในการจัดกลุ่มสิ่งทดลอง คือ แต่ละระดับของแต่ละปัจจัยต้องมารวมกัน ยกตัวอย่างเช่น กรณีศึกษาปัจจัย A และ ปัจจัย B ที่แต่ละปัจจัยมี 2 ระดับ (ระดับต่ำ และ ระดับสูง) สามารถจัดกลุ่มสิ่งทดลองได้ ดังนี้

วิธีที่ 1 ใช้ตารางจัดกลุ่มสิ่งทดลอง

		ระดับของปัจจัย A	
		a_1	a_2
ระดับของปัจจัย B	b_1	a_1b_1	a_2b_1
	b_2	a_1b_2	a_2b_2

วิธีที่ 2 ใช้แผนภูมิแกว่งปลา



จากกรณีข้างต้น จะสามารถจัดกลุ่มสิ่งทดลองได้ 4 สิ่งทดลอง (Treatment) ได้แก่ a_1b_1 , a_1b_2 , a_2b_1 และ a_2b_2 ในหนังสือเล่มนี้จะอธิบายตัวอย่างการจัดกลุ่มสิ่งทดลองที่มี 2 ปัจจัย และ 3 ปัจจัย ดังนี้

10.2.1 การจัดกลุ่มสิ่งทดลองแบบแฟกทอเรียลชนิด 2 ปัจจัย แต่ละปัจจัยมี 2 ระดับ

การจัดกลุ่มสิ่งทดลองที่มี 2 ปัจจัย และแต่ละปัจจัยมี 2 ระดับ เรียกว่าการจัดสิ่งทดลองแบบ 2×2 หรือ 2^2 แฟกทอเรียล จัดเป็นการทดลองที่ง่ายที่สุดซึ่งมักใช้ประกอบการวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) และการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) กล่าวคือ หากไม่มีการบล็อกในการทำซ้ำ หรือเป็นแบบสุ่มสมบูรณ์ จัดเป็น 2^2 แฟกทอเรียล ใน CRD (2^2 Factorial in CRD) หากมีการบล็อกโดยใช้ บล็อกเป็นซ้ำ จัดเป็น 2^2 แฟกทอเรียล ใน RCBD (2^2 Factorial in RCBD)

ขั้นตอนการจัดกลุ่มสิ่งทดลองแบบ 2^2 แฟกทอเรียล

- (1) กำหนดชื่อปัจจัยทั้ง 2 ปัจจัย ได้แก่ ปัจจัย A และ ปัจจัย B
- (2) กำหนดระดับของปัจจัย ในที่นี้แต่ละปัจจัยมี 2 ระดับ คือ ระดับต่ำ และ ระดับสูง สามารถใช้สัญลักษณ์ในการกำหนดระดับของปัจจัย ได้ 3 แบบ ดังนี้

แบบที่ 1 กำหนดระดับของแต่ละปัจจัยโดยใช้อักษรภาษาอังกฤษตัวพิมพ์เล็กที่มีเลขกำกับ ได้แก่

a_1 แทน ปัจจัย A ที่ระดับต่ำ a_2 แทน ปัจจัย A ที่ระดับสูง
 b_1 แทน ปัจจัย B ที่ระดับต่ำ b_2 แทน ปัจจัย B ที่ระดับสูง

การเขียนสัญลักษณ์ของสิ่งทดลองที่เกิดจากการรวมกันของปัจจัย (Treatment combination) คือ การนำอักษรภาษาอังกฤษตัวพิมพ์เล็กของทั้งสองปัจจัยมาไว้ด้วยกัน เช่น สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A ที่ระดับต่ำและปัจจัย B ที่ระดับสูง จะใช้สัญลักษณ์ a_1b_2

แบบที่ 2 กำหนดสิ่งทดลองที่เกิดจากการรวมกันของปัจจัย โดยใช้ลำดับมาตรฐาน ดังนี้

- (1) แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A และ B ที่ระดับต่ำ
- a แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A ระดับสูง และ B ระดับต่ำ
- b แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย B ระดับสูง และ A ระดับต่ำ
- ab แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A และ B ที่ระดับสูง

แบบที่ 3 กำหนดระดับของแต่ละปัจจัยโดยใช้เครื่องหมายบวกและลบ ดังนี้ ระดับต่ำ (ใช้เครื่องหมาย -) และระดับสูง (ใช้เครื่องหมาย +)

- (3) ทำตารางการจัดกลุ่มสิ่งทดลอง (Treatment combination) ในที่นี้จำนวนสิ่งทดลอง เท่ากับ 4 สิ่งทดลอง (คำนวณได้จากผลคูณของระดับแต่ละปัจจัย เท่ากับ 2×2) การกำหนดรายละเอียดระดับของปัจจัย ในแต่ละสิ่งทดลองแสดงได้หลายรูปแบบดังตารางที่ 10.1

ตารางที่ 10.1 เปรียบเทียบรูปแบบการใช้สัญลักษณ์แสดงการจัดกลุ่มสิ่งทดลอง (Treatment combination) แบบแฟกทอเรียลที่มี 2 ปัจจัย ปัจจัยละ 2 ระดับ (2^2 แฟกทอเรียล)

สิ่งทดลองที่	แบบที่ 1 ใช้อักษรภาษาอังกฤษ	แบบที่ 2 ใช้ลำดับมาตรฐาน	แบบที่ 3 ใช้เครื่องหมายบวกและลบ	
			ระดับของปัจจัย A	ระดับของปัจจัย B
1	a_1b_1	(1)	-	-
2	a_2b_1	a	+	-
3	a_1b_2	b	-	+
4	a_2b_2	ab	+	+

จากตารางที่ 10.1 เมื่อพิจารณารูปแบบที่ 2 การใช้ลำดับมาตรฐาน (Standard order) ในการแสดงการจัดกลุ่มสิ่งทดลองนั้น ให้สังเกต สิ่งทดลองแรกที่ใช้สัญลักษณ์ (1) หมายถึง สิ่งทดลองนั้นใช้ปัจจัย A และ B ที่ระดับต่ำทั้งคู่ ส่วนสิ่งทดลองที่ใช้สัญลักษณ์ a หมายถึง สิ่งทดลองที่มี a เด่น นั่นคือ สิ่งทดลองที่ใช้ปัจจัย A ระดับสูงและปัจจัย B ระดับต่ำ ส่วนสิ่งทดลองที่ใช้สัญลักษณ์ ab หมายถึง สิ่งทดลองนั้นใช้ปัจจัย A และ B ที่ระดับสูงทั้งคู่

10.2.2 การจัดกลุ่มสิ่งทดลองแบบแฟกทอเรียลชนิด 3 ปัจจัย แต่ละปัจจัยมี 2 ระดับ

กรณีที่น่าสนใจต้องการศึกษา 3 ปัจจัย และแต่ละปัจจัยมี 2 ระดับ เรียกว่า การจัดสิ่งทดลองแบบ $2 \times 2 \times 2$ หรือ 2^3 แฟกทอเรียล ทั้งนี้การศึกษา n ปัจจัย ที่แต่ละปัจจัยมีระดับที่เท่ากัน M ระดับ เรียกว่า **การจัดสิ่งทดลองแบบ M^n แฟกทอเรียล** โดยการจัดกลุ่มสิ่งทดลองจะได้จำนวนสิ่งทดลอง (Treatment) เท่ากับ M^n สิ่งทดลอง (n = จำนวนปัจจัย และ M = จำนวนระดับของปัจจัย) ยกตัวอย่างเช่น

2^2 หมายถึง Factorial ที่มีการศึกษา 2 ปัจจัย ปัจจัยละ 2 ระดับ จำนวนสิ่งทดลอง เท่ากับ 4 สิ่งทดลอง
 2^3 หมายถึง Factorial ที่มีการศึกษา 3 ปัจจัย ปัจจัยละ 2 ระดับ จำนวนสิ่งทดลอง เท่ากับ 8 สิ่งทดลอง
 3^2 หมายถึง Factorial ที่มีการศึกษา 2 ปัจจัย ปัจจัยละ 3 ระดับ จำนวนสิ่งทดลอง เท่ากับ 9 สิ่งทดลอง

*ข้อสังเกต การคำนวณจำนวนสิ่งทดลองโดยใช้สูตร M^n นี้ใช้ได้กรณีแต่ละปัจจัยมีจำนวนระดับเท่ากัน

ขั้นตอนการจัดกลุ่มสิ่งทดลองแบบ 2^3 แฟกทอเรียล

- (1) กำหนดชื่อปัจจัยทั้ง 3 ปัจจัย ได้แก่ ปัจจัย A, ปัจจัย B และ ปัจจัย C
- (2) กำหนดระดับของปัจจัย ในที่นี้แต่ละปัจจัยมี 2 ระดับ คือ ระดับต่ำ และ ระดับสูง สามารถใช้สัญลักษณ์ในการกำหนดระดับของปัจจัย ได้ 3 แบบ ดังนี้

แบบที่ 1 กำหนดระดับของแต่ละปัจจัยโดยใช้อักษรภาษาอังกฤษตัวพิมพ์เล็กที่มีเลขกำกับ ได้แก่

a_1 แทน ปัจจัย A ที่ระดับต่ำ	a_2 แทน ปัจจัย A ที่ระดับสูง
b_1 แทน ปัจจัย B ที่ระดับต่ำ	b_2 แทน ปัจจัย B ที่ระดับสูง
c_1 แทน ปัจจัย C ที่ระดับต่ำ	c_2 แทน ปัจจัย C ที่ระดับสูง

การเขียนสัญลักษณ์ของสิ่งทดลองที่เกิดจากการรวมกันของปัจจัย (Treatment combination) คือ การนำอักษรภาษาอังกฤษตัวพิมพ์เล็กของทั้ง 3 ปัจจัยมาไว้ด้วยกัน เช่น สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A ที่ระดับต่ำ, ปัจจัย B ที่ระดับสูง และปัจจัย C ที่ระดับสูง จะใช้สัญลักษณ์ $a_1b_2c_2$

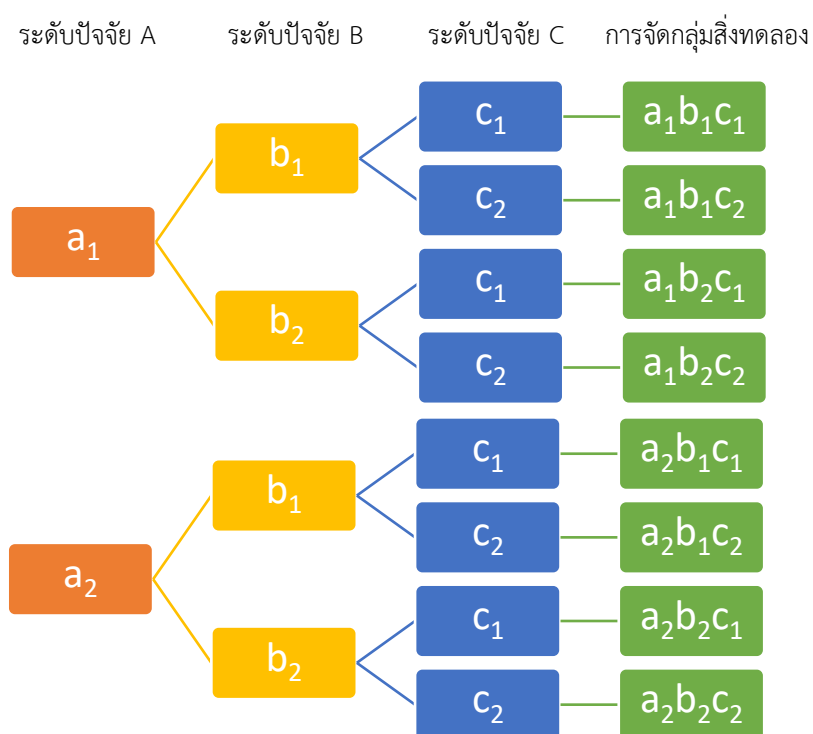
แบบที่ 2 กำหนดสิ่งทดลองที่เกิดจากการรวมกันของปัจจัย โดยใช้ลำดับมาตรฐาน ดังนี้

- (1) แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A, B และ C ที่ระดับต่ำ
- a แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A ระดับสูง, B และ C ระดับต่ำ
- b แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย B ระดับสูง, A และ C ระดับต่ำ
- ab แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A และ B ที่ระดับสูง, C ที่ระดับต่ำ
- c แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย C ระดับสูง, A และ B ระดับต่ำ
- ac แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A และ C ที่ระดับสูง, B ที่ระดับต่ำ
- bc แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย B และ C ที่ระดับสูง, A ที่ระดับต่ำ
- abc แทน สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A, B และ C ที่ระดับสูง

แบบที่ 3 กำหนดระดับของแต่ละปัจจัยโดยใช้เครื่องหมายบวกและลบ ดังนี้ ระดับต่ำ (ใช้เครื่องหมาย -) และระดับสูง (ใช้เครื่องหมาย +)

- (3) ทำตารางการจัดกลุ่มสิ่งทดลอง (Treatment combination) ในที่นี้จำนวนสิ่งทดลอง เท่ากับ 8 สิ่งทดลอง (คำนวณได้จากผลคูณของระดับแต่ละปัจจัย เท่ากับ $2 \times 2 \times 2$)

- (4) ใช้แผนภูมิแก๊งปลาจัดกลุ่มสิ่งทดลองแบบ 2^3 แฟกทอเรียล ดังภาพที่ 10.1 การกำหนดรายละเอียดระดับของปัจจัยในแต่ละสิ่งทดลองแสดงได้หลายรูปแบบดังตารางที่ 10.2



ภาพที่ 10.1 แผนภูมิกิ่งปลาแสดงการจัดกลุ่มสิ่งทดลองแบบ 2^3 แฟกทอเรียล
(ปัจจัย A, ปัจจัย B และ ปัจจัย C มี 2 ระดับเท่ากัน)

ตารางที่ 10.2 เปรียบเทียบรูปแบบการใช้สัญลักษณ์แสดงการจัดกลุ่มสิ่งทดลอง (Treatment combination) แบบแฟกทอเรียลที่มี 3 ปัจจัย ปัจจัยละ 2 ระดับ (2^3 แฟกทอเรียล)

สิ่งทดลองที่	แบบที่ 1 ใช้อักษร ภาษาอังกฤษ	แบบที่ 2 ใช้ลำดับ มาตรฐาน	แบบที่ 3 ใช้เครื่องหมายบวกและลบ		
			ระดับของ ปัจจัย A	ระดับของ ปัจจัย B	ระดับของ ปัจจัย C
1	$a_1b_1c_1$	(1)	-	-	-
2	$a_2b_1c_1$	a	+	-	-
3	$a_1b_2c_1$	b	-	+	-
4	$a_2b_2c_1$	ab	+	+	-
5	$a_1b_1c_2$	c	-	-	+
6	$a_2b_1c_2$	ac	+	-	+
7	$a_1b_2c_2$	bc	-	+	+
8	$a_2b_2c_2$	abc	+	+	+

10.2.3 การจัดกลุ่มสิ่งทดลองแบบแฟกทอเรียลชนิด 2 ปัจจัย แต่ละปัจจัยมีระดับไม่เท่ากัน

กรณีที่นักวิจัยต้องการศึกษาปัจจัย 2 ปัจจัยที่มีระดับของปัจจัยไม่เท่ากัน เช่น ต้องการศึกษากิจกรรม A ที่มี a ระดับ และปัจจัย B ที่มี b ระดับ เรียกว่า **การจัดสิ่งทดลองแบบ $a \times b$ แฟกทอเรียล** และเมื่อนำระดับของทุกปัจจัยมาจัดกลุ่มสิ่งทดลองจะได้จำนวนที่เป็นไปได้ทั้งหมด เท่ากับ **ab สิ่งทดลอง** และถ้าแต่ละสิ่งทดลองมีจำนวนซ้ำเท่ากับ n จะได้จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ **abn หน่วยทดลอง**

ยกตัวอย่างเช่น นักวิจัยต้องการศึกษาปัจจัย A ที่มี 2 ระดับ และปัจจัย B ที่มี 3 ระดับ จะเรียกว่าการจัดสิ่งทดลองแบบ 2×3 แฟกทอเรียล โดยมีจำนวนสิ่งทดลอง เท่ากับ 6 สิ่งทดลอง

ขั้นตอนการจัดกลุ่มสิ่งทดลองแบบ 2×3 แฟกทอเรียล

- (1) กำหนดชื่อปัจจัยทั้ง 2 ปัจจัย ได้แก่ ปัจจัย A และ ปัจจัย B
- (2) กำหนดระดับของปัจจัย ในที่นี้ปัจจัย A มี 2 ระดับ คือ ระดับต่ำ และ ระดับสูง และ ปัจจัย B มี 3 ระดับ คือ ระดับต่ำ, ระดับกลาง และระดับสูง สามารถใช้สัญลักษณ์ในการกำหนดระดับของปัจจัย ได้ 2 แบบ ดังนี้

แบบที่ 1 กำหนดระดับของแต่ละปัจจัยโดยใช้อักษรภาษาอังกฤษตัวพิมพ์เล็กที่มีเลขกำกับ ได้แก่

กำหนดปัจจัย A : a_1 แทน ปัจจัย A ที่ระดับต่ำ

a_2 แทน ปัจจัย A ที่ระดับสูง

กำหนดปัจจัย B : b_1 แทน ปัจจัย B ที่ระดับต่ำ

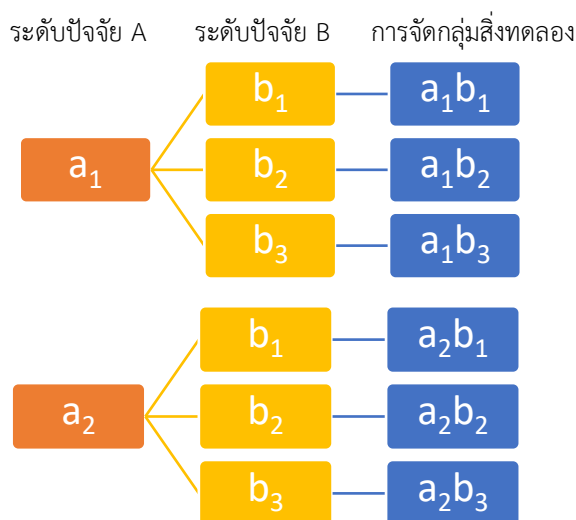
b_2 แทน ปัจจัย B ที่ระดับกลาง

b_3 แทน ปัจจัย B ที่ระดับสูง

การเขียนสัญลักษณ์ของสิ่งทดลองที่เกิดจากการรวมกันของปัจจัย (Treatment combination) คือ การนำอักษรภาษาอังกฤษตัวพิมพ์เล็กของทั้งสองปัจจัยมาไว้ด้วยกัน เช่น สิ่งทดลองที่เกิดจากการรวมกันของปัจจัย A ที่ระดับต่ำและปัจจัย B ที่ระดับกลาง จะใช้สัญลักษณ์ a_1b_2

แบบที่ 2 กำหนดระดับของแต่ละปัจจัยโดยใช้เครื่องหมายบวกและลบ ดังนี้ ระดับต่ำ (ใช้เครื่องหมาย -), ระดับกลาง (ใช้สัญลักษณ์ 0) และระดับสูง (ใช้เครื่องหมาย +)

- (3) ทำตารางการจัดกลุ่มสิ่งทดลอง ในที่นี้จำนวนสิ่งทดลอง เท่ากับ 6 สิ่งทดลอง (คำนวณได้จากผลคูณของระดับแต่ละปัจจัย เท่ากับ 2×3) ใช้แผนภูมิแก่งปลาจัดกลุ่มสิ่งทดลองแบบ 2×3 แฟกทอเรียล ได้ดังภาพที่ 10.2 การกำหนดรายละเอียดระดับของปัจจัยในแต่ละสิ่งทดลองแสดงดังตารางที่ 10.3



ภาพที่ 10.2 แผนภูมิผังปลาแสดงการจัดกลุ่มสิ่งทดลองแบบ 2×3 แฟกทอเรียล
(ปัจจัย A มี 2 ระดับ และ ปัจจัย B มี 3 ระดับ)

ตารางที่ 10.3 เปรียบเทียบรูปแบบการใช้สัญลักษณ์แสดงการจัดกลุ่มสิ่งทดลอง (Treatment combination) แบบ 2×3 แฟกทอเรียล

สิ่งทดลองที่	แบบที่ 1 ใช้อักษรภาษาอังกฤษ	แบบที่ 2 ใช้เครื่องหมายบวกและลบ	
		ระดับของปัจจัย A	ระดับของปัจจัย B
1	a_1b_1	-	-
2	a_1b_2	-	0
3	a_1b_3	-	+
4	a_2b_1	+	-
5	a_2b_2	+	0
6	a_2b_3	+	+

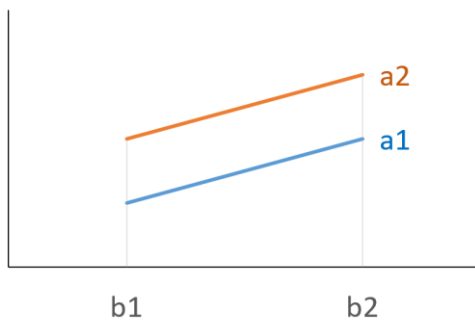
10.3 นิยามของอิทธิพลหลักและอิทธิพลร่วมระหว่างปัจจัย

10.3.1 อิทธิพลหลัก (Main effects) หมายถึง อิทธิพลของปัจจัยที่เมื่อระดับของปัจจัยเปลี่ยนไปจะทำให้ค่าสังเกตที่ได้จากการทดลองเปลี่ยนไปด้วย หรือกล่าวได้ว่า อิทธิพลของปัจจัยหลักที่ศึกษานั้นเอง สัญลักษณ์ที่ใช้แทนอิทธิพลหลัก คือ การใช้อักษรภาษาอังกฤษตัวพิมพ์ใหญ่ เช่น A, B, C, ...

10.3.2 อิทธิพลร่วมระหว่างปัจจัย หรือ ปฏิสัมพันธ์ระหว่างปัจจัย (Interaction effects) หมายถึง การเปลี่ยนแปลงในอิทธิพลของปัจจัยหนึ่งที่เกิดขึ้นเมื่อระดับของอีกปัจจัยหนึ่งเปลี่ยนแปลง การแสดง

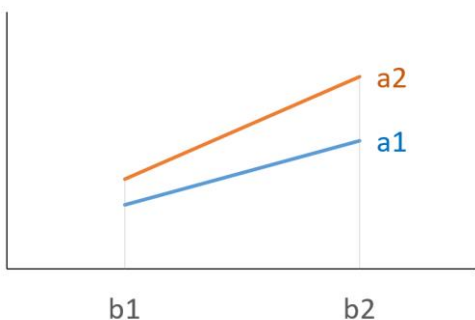
ปฏิสัมพันธ์ระหว่างปัจจัย A และ ปัจจัย B สามารถแสดงได้ดังภาพที่ 10.3-10.5 ทั้งนี้การมีหรือไม่มีอิทธิพลของปัจจัยหลักไม่สามารถบอกได้ว่ามีหรือไม่มีปฏิสัมพันธ์ของปัจจัยนั้น และเช่นเดียวกันการมีหรือไม่มีปฏิสัมพันธ์ระหว่างปัจจัยก็ไม่สามารถบอกได้ว่ามีหรือไม่มีอิทธิพลของปัจจัยหลักต่อค่าสังเกต แต่จะบอกเกี่ยวกับความสอดคล้องกันของอิทธิพลอย่างง่าย ๆ (Homogeneity) ความหมายของการทดสอบเมื่อมีปฏิสัมพันธ์ของปัจจัยหมายความว่าปัจจัยดังกล่าวไม่เป็นอิสระต่อกัน

ค่าสังเกต (Response scale)



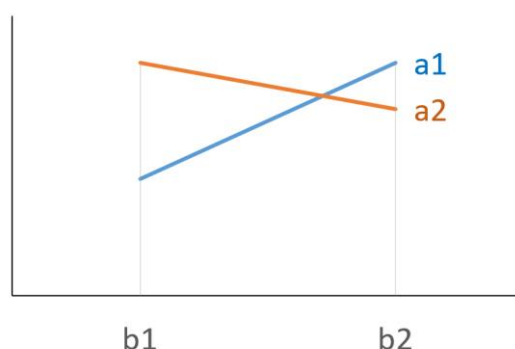
ภาพที่ 10.3 ไม่มีปฏิสัมพันธ์เกิดขึ้น

ค่าสังเกต (Response scale)



ภาพที่ 10.4 มีปฏิสัมพันธ์ของปัจจัยซึ่งเกิดขึ้นจากขนาดของการตอบสนองต่างกันเมื่อเปลี่ยนระดับของอีกปัจจัยหนึ่ง

ค่าสังเกต (Response scale)



ภาพที่ 10.5 มีปฏิสัมพันธ์ของปัจจัยซึ่งเกิดขึ้นจากความแตกต่างในทิศทางของการตอบสนอง

ในหนังสือเล่มนี้จะอธิบายการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย ในหัวข้อ 10.4 และ การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย ในหัวข้อ 10.6

10.4 การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย

การวางแผนการทดลองแบบ Factorial in CRD ชนิดที่มี 2 ปัจจัย จัดเป็นการทดลองที่ง่ายที่สุด เป็นการจัดสิ่งทดลองที่มี 2 ปัจจัย ได้แก่ ปัจจัย A มี a ระดับ และปัจจัย B มี b ระดับ เมื่อนำระดับมาจัดกลุ่ม สิ่งทดลองจะได้จำนวนสิ่งทดลองที่เป็นไปได้ เท่ากับ **ab** สิ่งทดลอง (ระดับของปัจจัย A × ระดับของปัจจัย B) และถ้าแต่ละสิ่งทดลองมีจำนวนซ้ำเท่ากับ **n** จะได้จำนวนหน่วยทดลองทั้งหมดเท่ากับ **abn** หน่วยทดลอง กรณีที่ปัจจัยเป็นอิทธิพลที่จะเขียนตัวแบบของแผนการทดลองแบบ Factorial in CRD ชนิดที่มี 2 ปัจจัย ได้ดังนี้

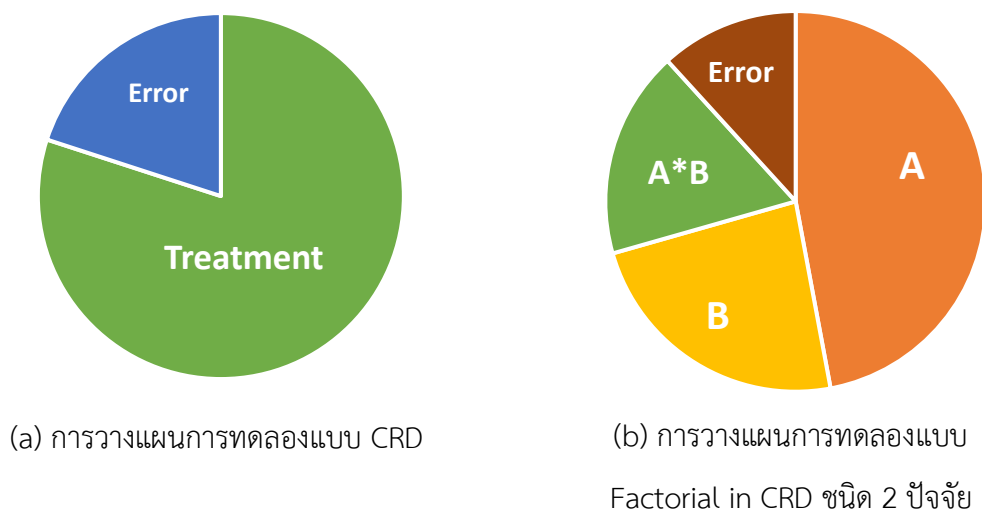
$$y_{ijk} = \mu + \tau_i + \beta_j + (\tau\beta)_{ij} + \varepsilon_{ijk} \quad \begin{matrix} i = 1, 2, \dots, a \\ j = 1, 2, \dots, b \\ k = 1, 2, \dots, n \end{matrix}$$

โดยที่ y_{ijk} คือ ค่าสังเกตที่ k ของปัจจัย A ที่ระดับ i และปัจจัย B ที่ระดับ j
 μ คือ ค่าเฉลี่ยของข้อมูลทั้งหมด
 τ_i คือ อิทธิพล (Effect) จากปัจจัย A ที่ระดับ i
 β_j คือ อิทธิพล (Effect) จากปัจจัย B ที่ระดับ j
 $(\tau\beta)_{ij}$ คือ อิทธิพล (Effect) ปฏิสัมพันธ์ระหว่าง τ_i และ β_j
 ε_{ijk} คือ Experimental error หรือ Random error

10.4.1 ความแปรปรวนของการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย

โดยทั่วไปแล้วการวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) เป็นการวางแผนการทดลองที่มี 1 ปัจจัย ดังนั้นจึงมีแหล่งของความแปรปรวน 2 แหล่ง คือ ความแปรปรวนที่อธิบายได้จากสิ่งทดลอง (Treatment) และ ความแปรปรวนที่อธิบายไม่ได้จากความคลาดเคลื่อนของการทดลอง (Experimental error) ดังภาพที่ 10.6(a) แต่ในกรณีที่วางแผนการทดลองแบบ Factorial in CRD และมีปัจจัยที่ต้องการศึกษา 2 ปัจจัย (ปัจจัย A และปัจจัย B) นั้นหมายความว่า สิ่งทดลองเกิดจากการจัดกลุ่มของระดับสองปัจจัยร่วมกัน ดังนั้นความแปรปรวนของสิ่งทดลองจึงประกอบด้วย ความแปรปรวนที่เกิดจากปัจจัย A, ความแปรปรวนที่เกิดจากปัจจัย B และความแปรปรวนที่เกิดจากปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B

กล่าวโดยสรุป คือ แหล่งของความแปรปรวนในการวางแผนการทดลองแบบ Factorial in CRD ชนิด ที่มี 2 ปัจจัย จะประกอบด้วย 4 แหล่ง แบ่งเป็น ความแปรปรวนที่อธิบายได้ 3 แหล่ง (จากปัจจัย A, ปัจจัย B และ ปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B) และ ความแปรปรวนที่อธิบายไม่ได้ 1 แหล่งจากความคลาดเคลื่อนของการทดลอง (Experimental error) ดังภาพที่ 10.6(b) ดังนั้นการวิเคราะห์ความแปรปรวนจะเป็นแบบสองทาง (Two-Way ANOVA) ซึ่งสามารถแสดงผลตารางการวิเคราะห์ความแปรปรวนได้ดังตารางที่ 10.4



ภาพที่ 10.6 เปรียบเทียบการแบ่งสัดส่วนความแปรปรวนของ (a) การวางแผนการทดลองแบบ CRD และ (b) การวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย

ตารางที่ 10.4 การวิเคราะห์ความแปรปรวนสำหรับการวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย โดยมีตัวแบบอิทธิพลแบบคงที่ (Fixed effect model)

แหล่งความแปรปรวน (Source of Variation)	องศาความเป็นอิสระ (Degree of Freedom)	ผลบวกกำลังสอง (Sum of Squares)	ค่ากำลังสองเฉลี่ย (Mean Squares)	F
ปัจจัย A	$a - 1$	SSA	$MSA = \frac{SSA}{a - 1}$	$F = \frac{MSA}{MSE}$
ปัจจัย B	$b - 1$	SSB	$MSB = \frac{SSB}{b - 1}$	$F = \frac{MSB}{MSE}$
ปฏิสัมพันธ์ AB	$(a - 1)(b - 1)$	SSAB	$MSAB = \frac{SSAB}{(a - 1)(b - 1)}$	$F = \frac{MSAB}{MSE}$
ความคลาดเคลื่อน	$ab(n - 1)$	SSE	$MSE = \frac{SSE}{ab(n - 1)}$	
ยอดรวม (Total)	$abn - 1$	SST		

a = จำนวนสิ่งทดลองหรือระดับของปัจจัย A

b = จำนวนสิ่งทดลองหรือระดับของปัจจัย B

10.4.2 การเขียนสมมติฐานทางสถิติสำหรับการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย

ในการทดสอบสมมติฐานสำหรับการวางแผนการทดลองแบบ Factorial in CRD ชนิดที่มี 2 ปัจจัย คือ ปัจจัย A และ ปัจจัย B จะให้ความสนใจในการทดสอบอิทธิพลของปัจจัยหลัก ได้แก่ อิทธิพลของปัจจัย A และอิทธิพลของปัจจัย B ที่มีต่อค่าสังเกต (ตัวแปรตาม) และอิทธิพลร่วมระหว่างระดับของปัจจัย A และปัจจัย B ที่มีต่อค่าสังเกต ดังนั้นการทดสอบสมมติฐานทางสถิติจะประกอบด้วย 3 ข้อ ดังนี้

สมมติฐานข้อ 1 การทดสอบอิทธิพลของปัจจัย A (Main effect A) ต่อตัวแปรตาม

H_0 : ปัจจัย A ไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : ปัจจัย A มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ H_0 : ค่าเฉลี่ยของตัวอย่างในทุกระดับของปัจจัย A ไม่แตกต่างกัน

H_1 : มีระดับของปัจจัย A อย่างน้อย 2 ระดับที่มีค่าเฉลี่ยแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\tau_i = 0$

H_1 : $\tau_i \neq 0$ อย่างน้อย 1 ค่า

สมมติฐานข้อ 2 การทดสอบอิทธิพลของปัจจัย B (Main effect B) ต่อตัวแปรตาม

H_0 : ปัจจัย B ไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : ปัจจัย B มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ H_0 : ค่าเฉลี่ยของตัวอย่างในทุกระดับของปัจจัย B ไม่แตกต่างกัน

H_1 : มีระดับของปัจจัย B อย่างน้อย 2 ระดับที่มีค่าเฉลี่ยแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\beta_j = 0$

H_1 : $\beta_j \neq 0$ อย่างน้อย 1 ค่า

สมมติฐานข้อ 3 การทดสอบอิทธิพลร่วมระหว่างปัจจัย A และปัจจัย B (Interaction effects A×B) ต่อตัวแปรตาม

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B ต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : มีปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B ต่อค่าเฉลี่ยของตัวแปรตาม

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $(\tau\beta)_{ij} = 0$

H_1 : $(\tau\beta)_{ij} \neq 0$ อย่างน้อย 1 ค่า

10.4.3 การทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Tests)

เมื่อวิเคราะห์ความแปรปรวนแล้วสรุปผลว่า ปัจจัยหลักมีอิทธิพลต่อตัวแปรตาม และ/หรือ มีอิทธิพลร่วมระหว่างปัจจัยหลักต่อตัวแปรตาม นักวิจัยส่วนใหญ่มักต้องการวิเคราะห์เปรียบเทียบความแตกต่างของค่าเฉลี่ยในแต่ละสิ่งทดลองที่เกิดจากการร่วมระดับของปัจจัย (Treatment combination) ในกรณีนี้ต้องทำการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ ยกตัวอย่างเช่น การวางแผนการทดลองแบบ 2^2 Factorial in CRD มีปัจจัยที่ศึกษา 2 ปัจจัย และแต่ละปัจจัยมี 2 ระดับ จะมีจำนวนสิ่งทดลองเท่ากับ 4 สิ่งทดลองให้ทำการเปรียบเทียบความแตกต่างของค่าเฉลี่ย

10.4.4 คำสั่ง SPSS ในการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in CRD ชนิด 2 ปัจจัย

คำสั่งที่ใช้ในการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in CRD ชนิดที่มี 2 ปัจจัย คือ การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) ด้วยคำสั่ง ดังนี้

Analyze ➤ General Linear Model ➤ Univariate...

10.4.4.1 ตัวอย่างข้อมูลการวางแผนการทดลองแบบ 2^2 Factorial in CRD

นักวิจัยต้องการศึกษาว่าอุณหภูมิและเวลาในการทอดมีผลต่อการอมน้ำมันของมันฝรั่งทอดหรือไม่ จึงได้วางแผนการทดลองแบบ 2^2 Factorial in CRD โดยกำหนดให้ปัจจัย A คือ อุณหภูมิในการทอด มี 2 ระดับ คือ 150 และ 160 °C และปัจจัย B คือ เวลาในการทอด มี 2 ระดับ คือ 3 และ 5 นาที ทำการสุ่มตัวอย่างมันฝรั่งทอดแต่ละสิ่งทดลองไปวิเคราะห์ปริมาณไขมันในห้องปฏิบัติการ โดยวิเคราะห์สิ่งทดลองละ 3 ซ้ำ ได้ข้อมูลดังตารางที่ 10.5

ตารางที่ 10.5 ปริมาณไขมันของมันฝรั่งทอดที่ใช้อุณหภูมิและเวลาในการทอดแตกต่างกัน

สิ่งทดลองที่	อุณหภูมิในการทอด (°C)	เวลาในการทอด (นาที)	ปริมาณไขมัน (%)		
			ซ้ำที่ 1	ซ้ำที่ 2	ซ้ำที่ 3
1	150	3	27.6	27.5	27.6
2	150	5	26.5	26.6	26.0
3	160	3	25.7	25.9	26.1
4	160	5	24.4	24.6	24.8

10.4.4.2 การเขียนสมมติฐาน

จากข้อมูลดังกล่าว สามารถสรุปรายละเอียดของข้อมูลได้ ดังนี้

- ปัจจัย A คือ อุณหภูมิในการทอด มี 2 ระดับ (150 และ 160 °C)
- ปัจจัย B คือ เวลาในการทอด มี 2 ระดับ (3 และ 5 นาที)
- จำนวนสิ่งทดลอง (Treatment) เท่ากับ 4 สิ่งทดลอง (ผลคูณของระดับแต่ละปัจจัย)
- จำนวนหน่วยทดลอง (Experimental unit) ในแต่ละสิ่งทดลอง เท่ากับ 3 หน่วย
- จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 12 หน่วย
- ค่าสังเกต หรือ ค่าตอบสนอง (ตัวแปรตาม) คือ ปริมาณไขมัน

นักวิจัยสามารถเขียนสมมติฐานทางสถิติที่ใช้ทดสอบได้ ดังนี้

(1) การทดสอบสมมติฐานอิทธิพลของอุณหภูมิในการทอด (Main effect A)

H_0 : อุณหภูมิในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

(2) การทดสอบสมมติฐานอิทธิพลของเวลาในการทอด (Main effect B)

H_0 : เวลาในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

(3) การทดสอบสมมติฐานอิทธิพลร่วมระหว่างอุณหภูมิในการทอดและเวลาในการทอด

(Interaction effects AxB)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมัน

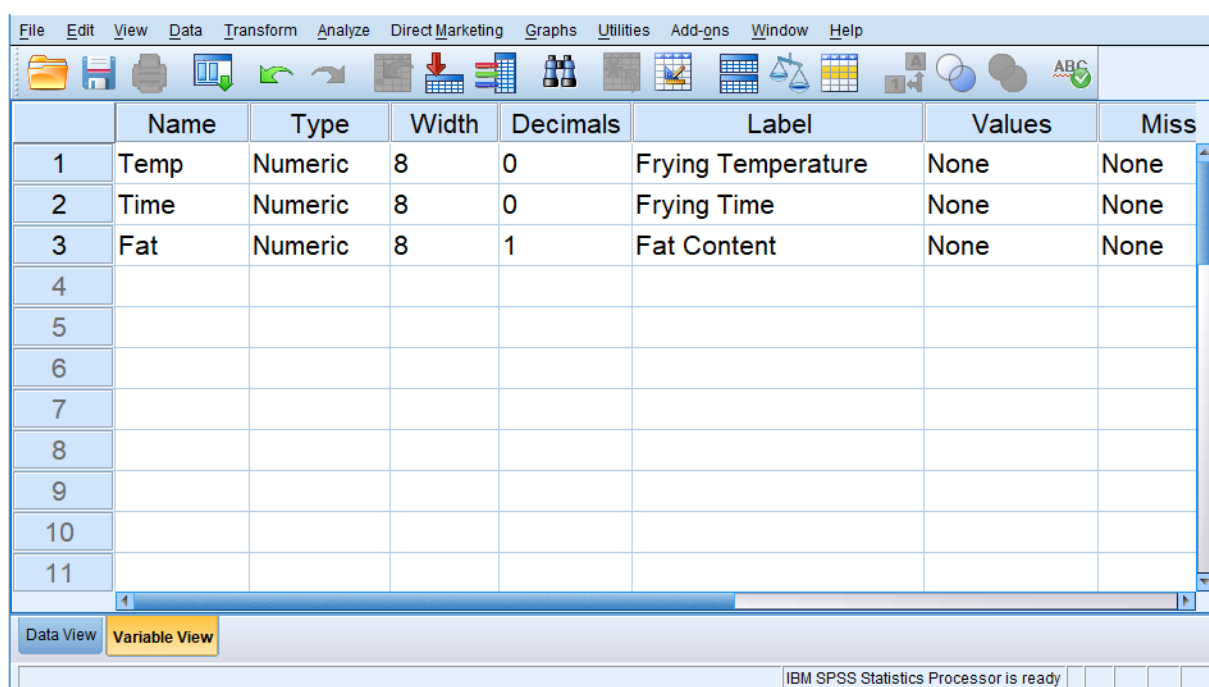
H_1 : มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมัน

10.4.4.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลการทดลองแบบ 2² Factorial in CRD

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลปริมาณไขมันของน้ำมันฟryingทอดแต่ละสิ่งทดลอง จะต้องใช้ตัวแปร var จำนวน 3 ตัวแปรสำหรับป้อนข้อมูล 3 ตัว ได้แก่ ระดับอุณหภูมิในการทอด, ระดับเวลาในการทอด และ ปริมาณไขมัน ดังนั้นกำหนดชื่อตัวแปรเป็น **Temp**, **Time** และ **Fat** ในหน้าต่าง Variable View ดังภาพที่ 10.7 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ระดับอุณหภูมิในการทอด	Temp	Frying Temperature	0	ไม่ต้องกำหนด
ระดับเวลาในการทอด	Time	Frying Time	0	ไม่ต้องกำหนด
ปริมาณไขมัน	Fat	Fat Content	1	ไม่ต้องกำหนด

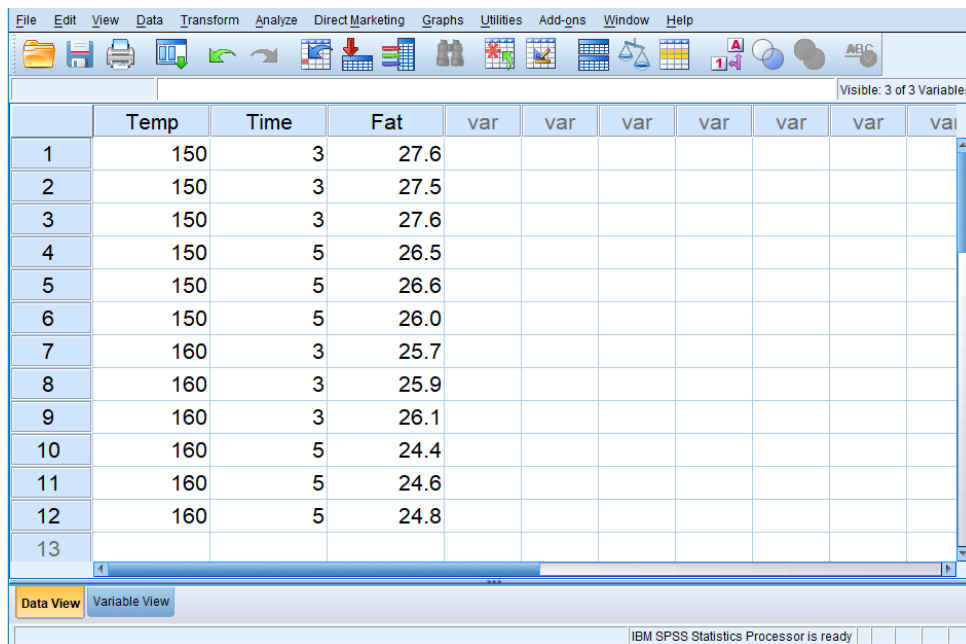
หมายเหตุ กรณีนักวิจัยจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



	Name	Type	Width	Decimals	Label	Values	Miss
1	Temp	Numeric	8	0	Frying Temperature	None	None
2	Time	Numeric	8	0	Frying Time	None	None
3	Fat	Numeric	8	1	Fat Content	None	None
4							
5							
6							
7							
8							
9							
10							
11							

ภาพที่ 10.7 กำหนดรายละเอียดตัวแปร Temp, Time และ Fat ใน Variable View

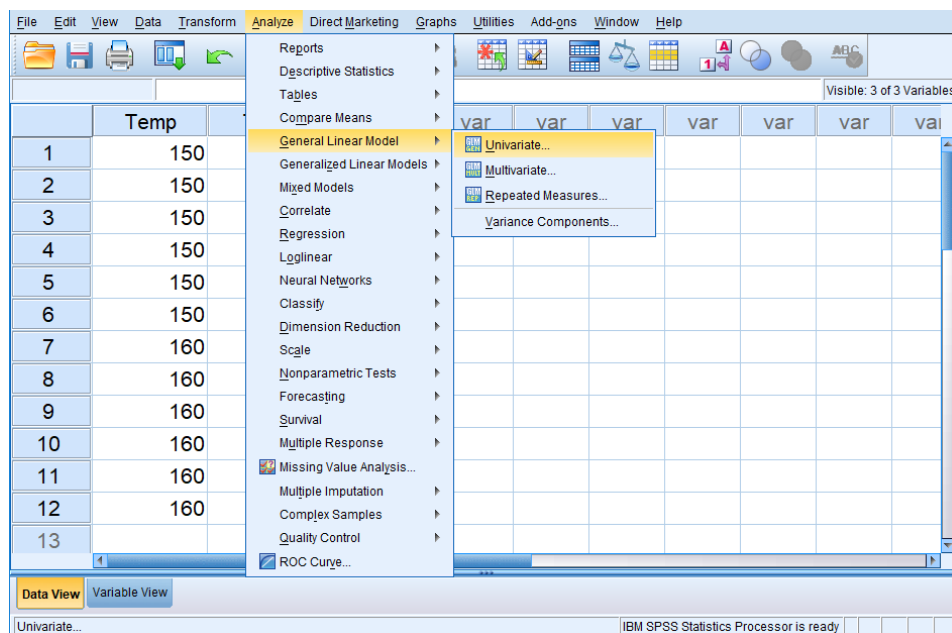
(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 10.8 โดยป้อนข้อมูลปริมาณไขมันเรียงตามระดับอุณหภูมิและเวลาในการทอด



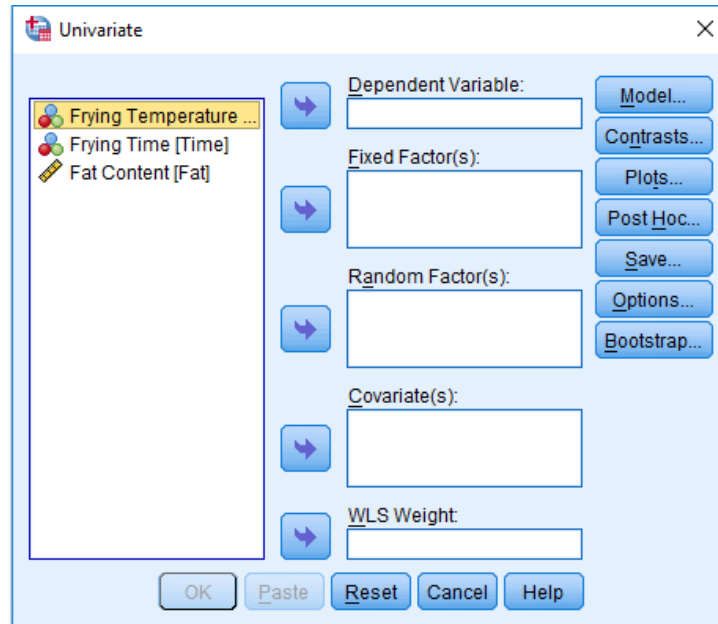
	Temp	Time	Fat	var	var	var	var	var	var	var
1	150	3	27.6							
2	150	3	27.5							
3	150	3	27.6							
4	150	5	26.5							
5	150	5	26.6							
6	150	5	26.0							
7	160	3	25.7							
8	160	3	25.9							
9	160	3	26.1							
10	160	5	24.4							
11	160	5	24.6							
12	160	5	24.8							
13										

ภาพที่ 10.8 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze ➤ General Linear Model ➤ Univariate...** ดังภาพที่ 10.9 จะปรากฏหน้าจอ ดังภาพที่ 10.10

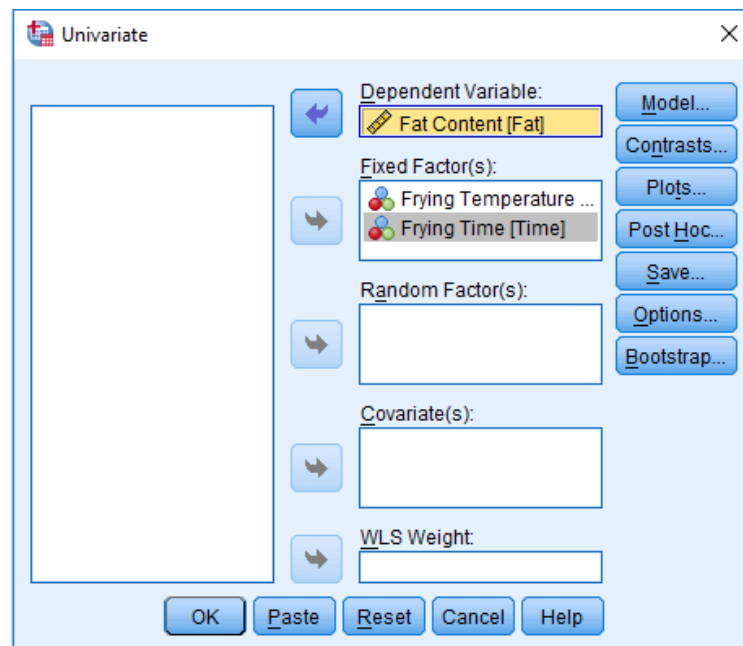


ภาพที่ 10.9 เมนูหน้าจอคำสั่ง Analyze ➤ General Linear Model ➤ Univariate...



ภาพที่ 10.10 หน้าจอคำสั่ง Univariate

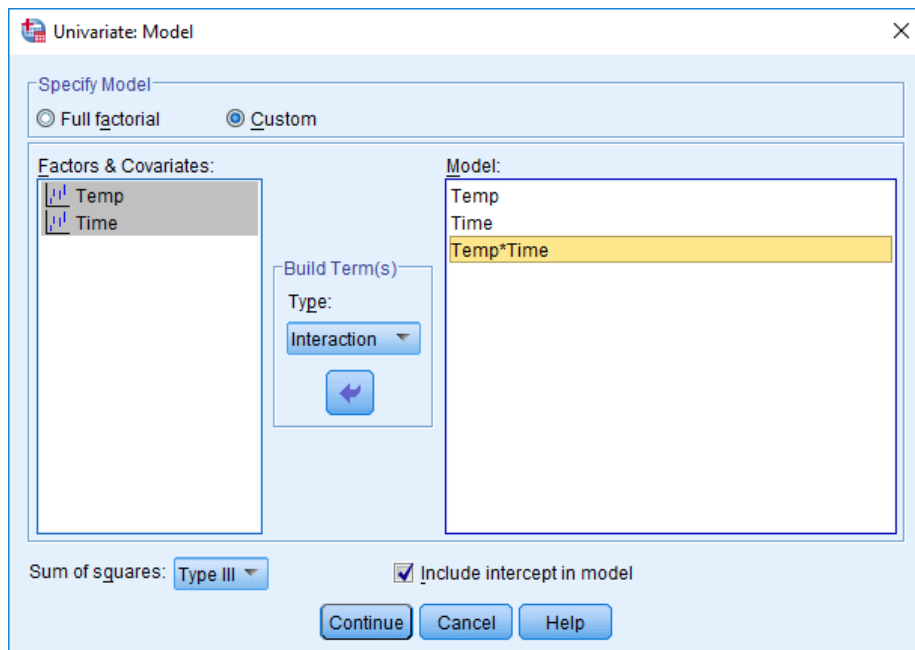
(4) เลือก ตัวแปร Fat Content[Fat] (ตัวแปรตาม) ใส่ในช่อง **Dependent Variable** และ ตัวแปร Frying Temperature[Temp] และ ตัวแปร Frying Time[Time] (ตัวแปรต้น) ใส่ในช่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.11



ภาพที่ 10.11 เลือกตัวแปร Fat Content[Fat] ใส่ในช่อง Dependent Variable และ เลือกตัวแปร Frying Temperature[Temp] และ ตัวแปร Frying Time[Time] ใส่ในช่อง Fixed Factor(s)

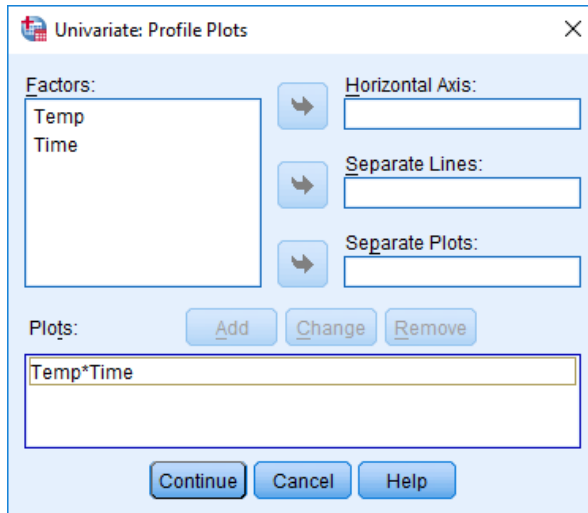
(5) เลือกคำสั่ง **Model...** จะได้หน้าจอตั้งภาพที่ 10.12 กำหนดค่าต่างๆ ดังนี้

- กำหนดรูปแบบการวิเคราะห์ (Model) ที่ตำแหน่ง Specify Model ให้คลิกเลือก **Custom** ทั้งนี้รูปแบบการวิเคราะห์ (Model) มี 2 รูปแบบ Full factorial คือ รูปแบบเต็มองค์ประกอบ และ Custom คือ การกำหนดรูปแบบขึ้นเอง
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลของปัจจัยหลัก (Main Effects) โดยกด Build Term(s) เลือก **Main Effects** แล้วเลือกตัวแปร **Temp** และ **Time** ในกล่องซ้ายมือ ใส่ในช่อง Model ที่อยู่ในกล่องขวามือ
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลร่วมระหว่างปัจจัยทั้ง 2 ปัจจัย (Interaction Effects) โดยกด Build Term(s) เลือก **Interaction** แล้วกดคีย์บอร์ด **Shift ↑** ค้างไว้ จากนั้นกดเลือกตัวแปร Temp และ Time (กดเลือกแบบจับคู่) ในกล่องซ้ายมือ ใส่ในช่อง Model จะได้ตัวแปร **Temp*Time** ในกล่องขวามือ
- เมื่อกำหนดค่าเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



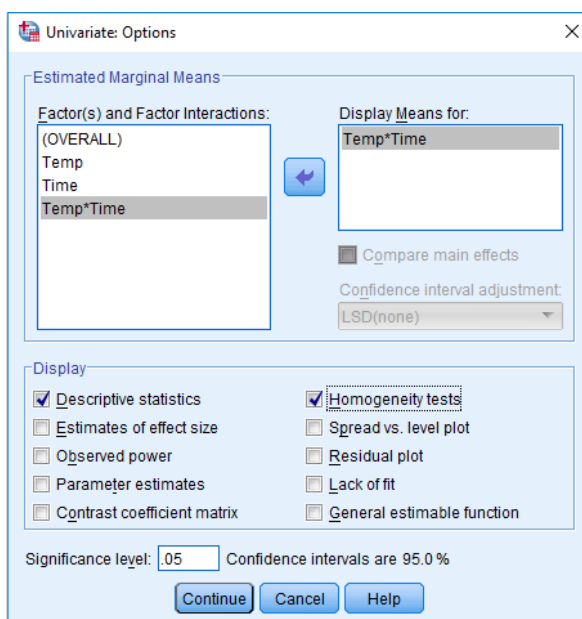
ภาพที่ 10.12 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลร่วมระหว่างทั้งสองปัจจัย จะได้หน้าจอตั้งภาพที่ 10.13 กดเลือกตัวแปร Temp ในช่อง Horizontal Axis (แกนนอน) และตัวแปร Time ในช่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Temp*Time** ในช่องด้านล่างเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 10.13 หน้าจอคำสั่ง Univariate : Profile Plots

(7) เลือกคำสั่ง **Options...** เพื่อให้แสดงค่าสถิติเบื้องต้นของแต่ละกลุ่ม จะได้หน้าจอตั้งภาพที่ 10.14 กดเลือก **ตัวแปร Temp*Time** ใส่ในช่อง Display Means for แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) เมื่อเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate และกด OK จะได้ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 10.15



ภาพที่ 10.14 หน้าจอคำสั่ง Univariate : Options

Univariate Analysis of Variance

Between-Subjects Factors

	N
Frying Temperature 150	6
160	6
Frying Time 3	6
5	6

Descriptive Statistics

Dependent Variable: Fat Content

Frying Temperature	Frying Time	Mean	Std. Deviation	N
150	3	27.567	.0577	3
	5	26.367	.3215	3
	Total	26.967	.6890	6
160	3	25.900	.2000	3
	5	24.600	.2000	3
	Total	25.250	.7342	6
Total	3	26.733	.9223	6
	5	25.483	.9968	6
	Total	26.108	1.1245	12

Levene's Test of Equality of Error Variances^a

Dependent Variable: Fat Content

F	df1	df2	Sig.
1.976	3	8	.196

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Temp + Time + Temp * Time

Tests of Between-Subjects Effects

Dependent Variable: Fat Content

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	13.536 ^a	3	4.512	96.685	.000
Intercept	8179.741	1	8179.741	175280.161	.000
Temp	8.841	1	8.841	189.446	.000
Time	4.688	1	4.688	100.446	.000
Temp * Time	.008	1	.008	.161	.699
Error	.373	8	.047		
Total	8193.650	12			
Corrected Total	13.909	11			

a. R Squared = .973 (Adjusted R Squared = .963)

Estimated Marginal Means

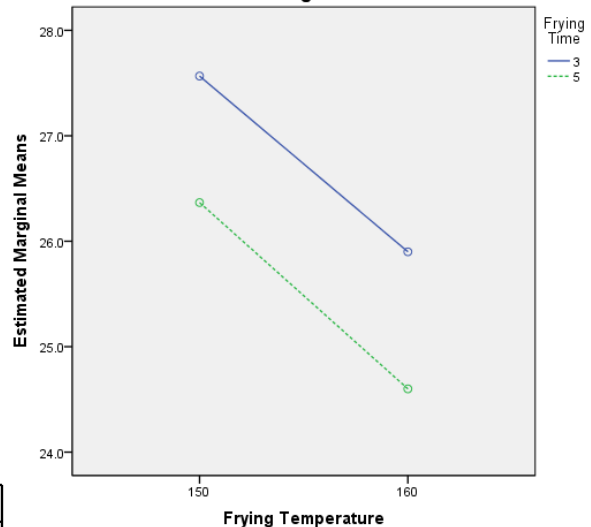
Frying Temperature * Frying Time

Dependent Variable: Fat Content

Frying Temperature	Frying Time	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
150	3	27.567	.125	27.279	27.854
	5	26.367	.125	26.079	26.654
160	3	25.900	.125	25.612	26.188
	5	24.600	.125	24.312	24.888

Profile Plots

Estimated Marginal Means of Fat Content



ภาพที่ 10.15 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนข้อมูลการทดลองแบบ 2^2 Factorial in CRD

10.4.4.4 การอ่านผลลัพธ์จากการวิเคราะห์ข้อมูลการทดลองแบบ 2^2 Factorial in CRD

จากผลการวิเคราะห์ในภาพที่ 10.15 แสดงผลลัพธ์ได้ 6 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Between-Subjects Factors เป็นส่วนที่แสดงรายละเอียดของข้อมูลที่ป้อน ดังนี้

Frying Temperature	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 150 และ 160 °C
Frying Time	คือ	ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 3 และ 5 นาที
N	คือ	จำนวนข้อมูลในแต่ละกลุ่มที่จำแนกโดยตัวแปร Frying Temperature และตัวแปร Frying Time ในที่นี้ แต่ละกลุ่มมีจำนวนข้อมูลเท่ากัน คือ 6 ค่า

ส่วนที่ 2 ตาราง Descriptive Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Frying Temperature	คือ	คือ ตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 150 และ 160 °C
Frying Time	คือ	คือ ตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 3 และ 5 นาที
Mean	คือ	คือ ค่าเฉลี่ยปริมาณไขมันของแต่ละกลุ่ม
Std. Deviation	คือ	คือ ค่าส่วนเบี่ยงเบนมาตรฐานของปริมาณไขมันที่แสดงการกระจายของข้อมูล
N	คือ	คือ จำนวนข้อมูลในแต่ละกลุ่มที่จำแนกด้วยตัวแปร Frying Temperature และ ตัวแปร Frying Time

ส่วนที่ 3 ตาราง Levene's Test of Equality of Error Variances ใช้พิจารณาการกระจายของข้อมูลหรือความแปรปรวนของข้อมูลแต่ละกลุ่ม อธิบายได้ ดังนี้

F, df1, df2	คือ	ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในกรณีนี้การทดสอบสมมติฐานเป็นแบบสองหาง (Two-tailed test) เขียนได้ดังนี้ H_0 : ข้อมูลทุกกลุ่มมีการกระจายตัวไม่แตกต่างกัน H_1 : มีข้อมูลอย่างน้อย 2 กลุ่มมีการกระจายตัวแตกต่างกัน ค่า Sig. มีค่าเท่ากับ 0.196 ซึ่งมากกว่าค่า 0.05 (ค่า α ที่นักวิจัยกำหนดไว้) ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ข้อมูลทุกกลุ่มมีการกระจายตัวไม่แตกต่างกันที่ ระดับนัยสำคัญ 0.05 (Equal variances assumed)

ส่วนที่ 4 ตาราง Tests of Between-Subjects Effects เป็นส่วนที่แสดงค่าสถิติต่างๆ ของตารางวิเคราะห์ความแปรปรวน เพื่อใช้ในการทดสอบสมมติฐานทางสถิติด้วยค่าสถิติ F หรือค่าความน่าจะเป็น Sig. เพื่อใช้ทดสอบสมมติฐาน อธิบายได้ ดังนี้

Type III Sum of Squares, df, Mean Square, F	คือ	ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในที่นี้จะพิจารณาค่า Sig. เพื่อทดสอบอิทธิพลของปัจจัยหลักทั้ง 2 ปัจจัย (ตัวแปร Temp และ ตัวแปร Time) และอิทธิพลร่วมระหว่าง 2 ปัจจัยนั้น (ตัวแปร Temp*Time) ตามสมมติฐานที่เขียนไว้

ส่วนที่ 5 ตาราง Estimated Marginal Means เป็นส่วนที่แสดงค่าเฉลี่ยโดยประมาณของตัวแปรที่นำมาทดสอบ เนื่องจากตอนวิเคราะห์เลือกตัวแปร Temp*Time ดังนั้นผลที่แสดงในตารางนี้จะเป็นค่าเฉลี่ยโดยประมาณของตัวแปร Frying Temperature แต่ละกลุ่มที่แยกตามระดับตัวแปร Frying Time จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Mean	คือ	ค่าเฉลี่ยปริมาณไขมันโดยประมาณของแต่ละกลุ่มที่จำแนกด้วยตัวแปร Frying Temperature และตัวแปร Frying Time
Std. Error	คือ	ค่าความคลาดเคลื่อนมาตรฐานในแต่ละกลุ่ม
95% Confidence Interval	คือ	ค่าที่แสดงขอบเขตบนล่างของค่าเฉลี่ยในแต่ละกลุ่มในช่วงความเชื่อมั่น 95%

ส่วนที่ 6 Profile Plots เป็นส่วนที่แสดงกราฟเส้นของตัวแปรที่ต้องการทดสอบ

ในที่นี้ตัวแปรที่ต้องการทดสอบ คือ ตัวแปร Fat กราฟที่แสดงใช้ค่าเฉลี่ยโดยประมาณ (จากตาราง Estimated Marginal Means) มาแสดงโดยจำแนกตามตัวแปรที่ได้เลือกในขั้นตอนการวิเคราะห์ด้วยคำสั่ง **Plots...** ในที่นี้ได้เลือกตัวแปร Temp เป็นตัวแปรแกนนอน จึงเห็นกราฟในแกนนอนมี 2 สเกล (แสดงอุณหภูมิในการทอดที่ 150 และ 160 °C) และได้เลือกตัวแปร Time แสดงเป็นเส้นกราฟ จึงเห็นในกราฟมี 2 เส้น (แสดงเวลาในการทอดที่ 3 และ 5 นาที) จากกราฟ พบว่า มีลักษณะเป็นเส้นขนานแสดงว่า 2 ปัจจัยนั้นไม่มีอิทธิพลร่วมกัน ทั้งนี้เนื่องจากการนำค่าโดยประมาณมาใช้จึงอาจเกิดความผิดพลาดได้ ในการทดสอบอิทธิพลร่วมระหว่างทั้ง 2 ปัจจัยจึงควรพิจารณาจากค่า Sig. ในตาราง Tests of Between-Subjects Effects ก่อนเป็นอันดับแรกและใช้กราฟ Profile Plots ในการพิจารณาแนวโน้มค่าเฉลี่ยระหว่างระดับต่างๆ ของปัจจัยทั้ง 2 ตัว

10.4.4.5 การสรุปผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2² Factorial in CRD

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่กำหนดไว้ ดังนี้

สมมติฐานข้อที่ 1 การทดสอบสมมติฐานอิทธิพลของอุณหภูมิในการทอด (Main effect A)

H_0 : อุณหภูมิในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.15 ค่า Sig. ของตัวแปร Temp มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 2 การทดสอบสมมติฐานอิทธิพลของเวลาในการทอด (Main effect B)

H_0 : เวลาในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.15 ค่า Sig. ของตัวแปร Time มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 3 การทดสอบสมมติฐานอิทธิพลร่วมระหว่างอุณหภูมิในการทอดและเวลาในการทอด (Interaction effects A×B)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.15 ค่า Sig. ของตัวแปร Temp*Time มีค่าเท่ากับ 0.699 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันที่ระดับนัยสำคัญ 0.05

10.4.4.6 การนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in CRD

การเขียนรายงานผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in CRD จะเป็นการเขียนนำเสนอผลลัพธ์เช่นเดียวกับการวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) ซึ่งสามารถนำเสนอได้หลายแบบ ในที่นี้จะนำเสนอในรูปแบบตาราง 2 รูปแบบดังตารางที่ 10.6 และ 10.7 ซึ่งนักวิจัยจะเลือกนำเสนอแบบใดนั้นขึ้นอยู่กับวัตถุประสงค์ที่ต้องการนำเสนอข้อมูล

ตารางที่ 10.6 ค่า P-value แสดงอิทธิพลของอุณหภูมิในการทอด อิทธิพลของเวลาในการทอด และอิทธิพลร่วมระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

อิทธิพลของปัจจัย	ค่า P-value
อุณหภูมิในการทอด	0.000*
เวลาในการทอด	0.000*
อุณหภูมิในการทอด × เวลาในการทอด	0.699

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.7 ค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ใช้อุณหภูมิและเวลาในการทอดต่างกัน

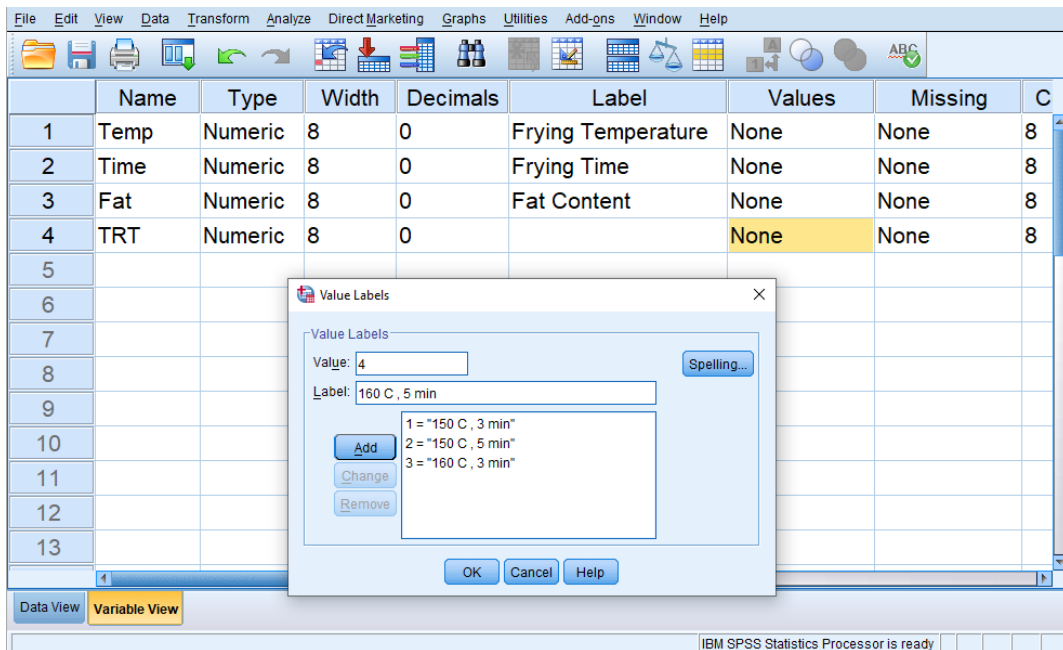
อุณหภูมิในการทอด (°C)	เวลาในการทอด (นาที)	ปริมาณไขมัน (%)
150	3	27.57 ± 0.06a
150	5	26.37 ± 0.32b
160	3	25.90 ± 0.20c
160	5	24.60 ± 0.20d
อิทธิพลของปัจจัย		ค่า P-value
อุณหภูมิในการทอด		0.000*
เวลาในการทอด		0.000*
อุณหภูมิในการทอด × เวลาในการทอด		0.699

^{a-d} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.7 ได้มาจากการกำหนดตัวแปร var ใหม่ชื่อ TRT ซึ่งเป็นตัวแปรที่แสดงถึงสิ่งทดลองที่เกิดจากปัจจัยร่วมระหว่างอุณหภูมิในการทอด (2 ระดับ) และ เวลาในการทอด (2 ระดับ) หรือเรียกว่า “Treatment Combination” ในตัวอย่างนี้จะได้จำนวนสิ่งทดลอง (TRT) เท่ากับ 4 สิ่งทดลอง (คำนวณจากผลคูณของระดับแต่ละปัจจัย) นำข้อมูลตัวแปร TRT และตัวแปร Fat ไปวิเคราะห์ด้วยคำสั่ง **Post Hoc...** เพื่อเปรียบเทียบค่าเฉลี่ยแบบจับคู่ทุกคู่ ทำเช่นเดียวกับการวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ CRD หรือ การวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) ตามขั้นตอนดังนี้

(1) กำหนดรายละเอียดและค่า Values ของตัวแปร TRT ในหน้าต่าง Variable View ดังภาพที่ 10.16 และป้อนข้อมูล TRT ในหน้าต่าง Data View ดังภาพที่ 10.17



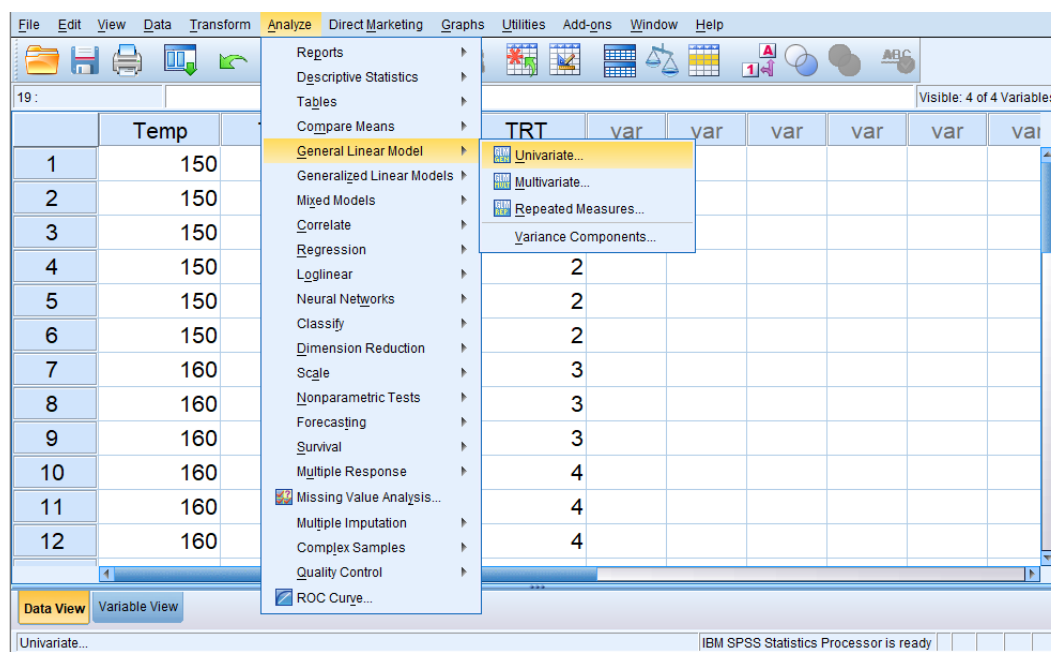
ภาพที่ 10.16 การกำหนดรายละเอียดตัวแปร TRT ในหน้าต่าง Variable View

(วิธีการกำหนดค่า Values ให้ตัวแปร สามารถอ่านเพิ่มเติมได้ในบทที่ 3)

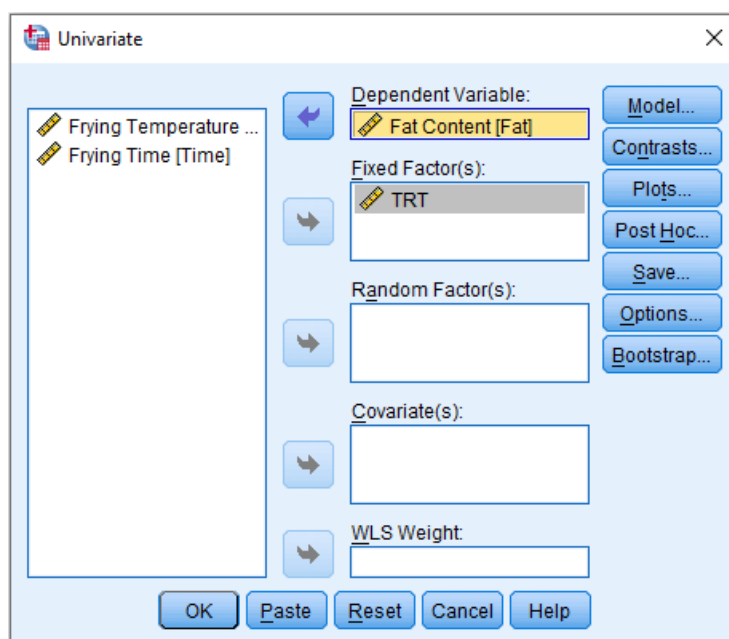
	Temp	Time	Fat	TRT	var	var	var	var	var	var
1	150	3	28	1						
2	150	3	28	1						
3	150	3	28	1						
4	150	5	27	2						
5	150	5	27	2						
6	150	5	26	2						
7	160	3	26	3						
8	160	3	26	3						
9	160	3	26	3						
10	160	5	24	4						
11	160	5	25	4						
12	160	5	25	4						

ภาพที่ 10.17 การป้อนข้อมูล TRT ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 10.18 และเลือกตัวแปร **Fat Content [Fat]** (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร **TRT** (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.19

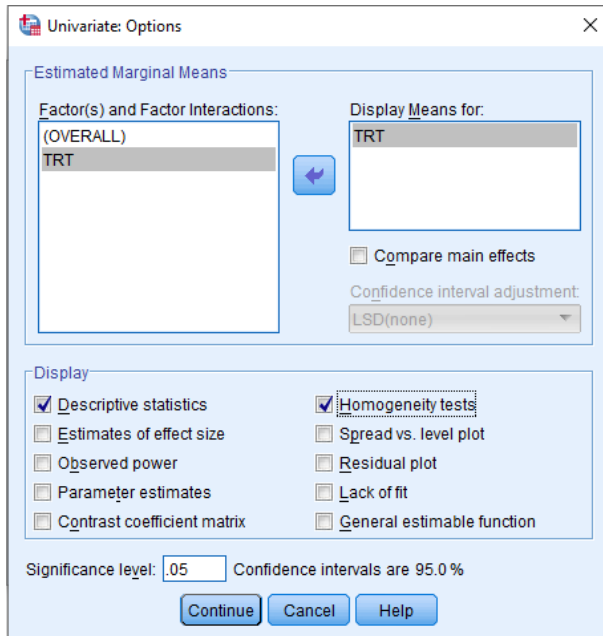


ภาพที่ 10.18 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



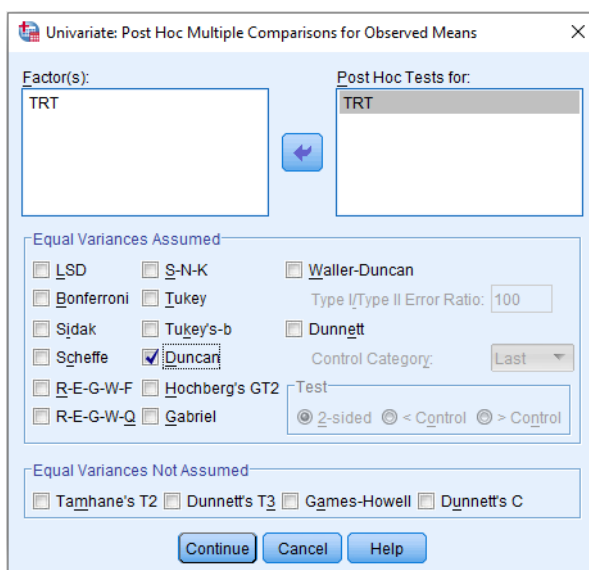
ภาพที่ 10.19 การกำหนดตัวแปรในเมนูคำสั่ง Univariate

(3) เลือกคำสั่ง **Options...** จะได้หน้าจอ ดังภาพที่ 10.20 กดเลือก **ตัวแปร TRT** ใส่ในกล่อง Display Means for ที่อยู่ด้านขวามือ และกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 10.20 การกำหนดตัวแปรในเมนูคำสั่ง Univariate : Options

(4) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอ ดังภาพที่ 10.21 ให้กดเลือก **ตัวแปร TRT** ใส่ในกล่อง Post Hoc Tests for: และกดเลือก **วิธี Duncan** ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 10.22



ภาพที่ 10.21 การกำหนดตัวแปรและวิธีการทดสอบในเมนูคำสั่ง Post Hoc...

Univariate Analysis of Variance

Between-Subjects Factors

	Value Label	N
TRT 1	150 C , 3 min	3
2	150 C , 5 min	3
3	160 C , 3 min	3
4	160 C , 5 min	3

Descriptive Statistics

Dependent Variable: Fat Content

TRT	Mean	Std. Deviation	N
150 C , 3 min	27.57	.058	3
150 C , 5 min	26.37	.321	3
160 C , 3 min	25.90	.200	3
160 C , 5 min	24.60	.200	3
Total	26.11	1.124	12

Levene's Test of Equality of Error Variances^a

Dependent Variable: Fat Content

F	df1	df2	Sig.
1.976	3	8	.196

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + TRT

Tests of Between-Subjects Effects

Dependent Variable: Fat Content

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	13.536 ^a	3	4.512	96.685	.000
Intercept	8179.741	1	8179.741	175280.161	.000
TRT	13.536	3	4.512	96.685	.000
Error	.373	8	.047		
Total	8193.650	12			
Corrected Total	13.909	11			

a. R Squared = .973 (Adjusted R Squared = .963)

Estimated Marginal Means

TRT

Dependent Variable: Fat Content

TRT	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
150 C , 3 min	27.567	.125	27.279	27.854
150 C , 5 min	26.367	.125	26.079	26.654
160 C , 3 min	25.900	.125	25.612	26.188
160 C , 5 min	24.600	.125	24.312	24.888

Post Hoc Tests

TRT

Homogeneous Subsets

Fat Content

Duncan^{a, b}

TRT	N	Subset			
		1	2	3	4
160 C , 5 min	3	24.60			
160 C , 3 min	3		25.90		
150 C , 5 min	3			26.37	
150 C , 3 min	3				27.57
Sig.		1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .047.

a. Uses Harmonic Mean Sample Size = 3.000.

b. Alpha = .05.

ภาพที่ 10.22 ผลลัพธ์การวิเคราะห์ความแปรปรวนของตัวแปร TRT และ ตัวแปร Fat

10.5 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 3x2 Factorial in CRD

10.5.1 ตัวอย่างข้อมูล

นักวิจัยได้ทำการสำรวจตลาดแล้วพบว่าผลิตภัณฑ์ที่มีส่วนผสมของทุเรียนกำลังเป็นที่ต้องการของตลาด ดังนั้นนักวิจัยจึงต้องการพัฒนาสูตรเค้กทุเรียน แต่เนื่องจากเนื้อทุเรียนสุกและน้ำตาลที่ใช้ในสูตรมีผลทำให้เค้กมีสีคล้ำอันเนื่องมาจากปฏิกิริยาสีน้ำตาล นักวิจัยจึงต้องการศึกษาอิทธิพลของปริมาณเนื้อทุเรียนและปริมาณน้ำตาลที่มีผลต่อค่าสี L^* ของเค้กทุเรียน โดยศึกษาปริมาณเนื้อทุเรียนสุกบด 3 ระดับ (30, 40 และ 50 กรัม) และปริมาณน้ำตาล 2 ระดับ (100 และ 150 กรัม) และควบคุมให้ส่วนผสมอื่นคงที่ วางแผนการทดลองแบบ 3x2 Factorial in CRD นำเค้กที่ได้แต่ละสิ่งทดลองไปวิเคราะห์ค่าสี L^* ด้วยเครื่อง Chroma meter จำนวน 3 ซ้ำ ได้ข้อมูลดังตารางที่ 10.8

ตารางที่ 10.8 ค่าสี L^* ของเค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียนและปริมาณน้ำตาลแตกต่างกัน

สิ่งทดลองที่	ปริมาณเนื้อ ทุเรียน (กรัม)	ปริมาณน้ำตาล (กรัม)	ค่าสี L^*		
			วัดซ้ำที่ 1	วัดซ้ำที่ 2	วัดซ้ำที่ 3
1	30	100	87.77	87.17	87.61
2	30	150	79.83	79.13	80.07
3	40	100	85.87	85.67	86.45
4	40	150	75.65	75.45	75.81
5	50	100	80.23	80.14	80.47
6	50	150	70.24	70.01	70.45

จากข้อมูลดังกล่าว พิจารณารายละเอียดของข้อมูลได้ ดังนี้

- ปัจจัย A คือ ปริมาณเนื้อทุเรียน มี 3 ระดับ (30, 40 และ 50 กรัม)
- ปัจจัย B คือ ปริมาณน้ำตาล มี 2 ระดับ (100 และ 150 กรัม)
- จำนวนสิ่งทดลอง (Treatment) เท่ากับ 6 สิ่งทดลอง (ผลคูณของระดับแต่ละปัจจัย)
- จำนวนหน่วยทดลอง (Experimental unit) ในแต่ละสิ่งทดลอง เท่ากับ 3 หน่วย
- จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 18 หน่วย
- ค่าสังเกต หรือ ค่าตอบสนอง (ตัวแปรตาม) คือ ค่าสี L^*

10.5.2 การเขียนสมมติฐานทางสถิติเพื่อวิเคราะห์ข้อมูลจากวางแผนการทดลองแบบ 3x2

Factorial in CRD

สามารถเขียนสมมติฐานสำหรับทดสอบทางสถิติได้ ดังนี้

สมมติฐานข้อที่ 1 ทดสอบอิทธิพลของระดับปริมาณเนื้อทุเรียน (Main effect A)

- H_0 : เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30, 40 และ 50 กรัม มีค่าเฉลี่ยสี L^* ไม่แตกต่างกัน
 H_1 : มีปริมาณเนื้อทุเรียน 2 ระดับที่ทำให้เค้กทุเรียนมีค่าเฉลี่ยสี L^* แตกต่างกัน

สมมติฐานข้อที่ 2 ทดสอบอิทธิพลของระดับปริมาณน้ำตาล (Main effect B)

- H_0 : เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยสี L^* ไม่แตกต่างกัน
 H_1 : เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยสี L^* แตกต่างกัน

สมมติฐานข้อที่ 3 ทดสอบอิทธิพลร่วมระหว่างระดับปริมาณเนื้อทุเรียนและปริมาณน้ำตาล

(Interaction effects AxB)

- H_0 : ไม่มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยสี L^* ของเค้ก
 H_1 : มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยสี L^* ของเค้ก

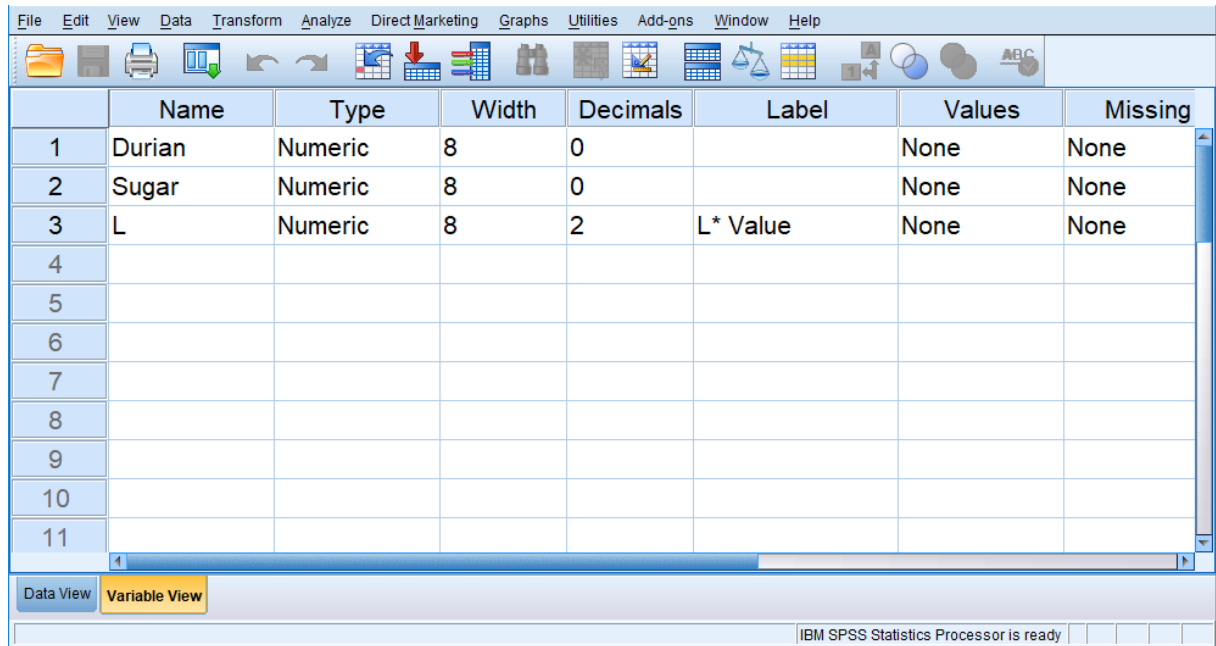
10.5.3 การใช้คำสั่งใน SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 3x2 Factorial in

CRD

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลค่าสี L^* ของเค้กทุเรียนแต่ละสิ่งทดลอง จะต้องใช้ตัวแปร var จำนวน 3 ตัวแปรสำหรับป้อนข้อมูล 3 ตัว ได้แก่ ปริมาณเนื้อทุเรียน, ปริมาณน้ำตาล และค่าสี L^* ดังนั้น กำหนดชื่อตัวแปรเป็น **Durian**, **Sugar** และ **L** ในหน้าต่าง Variable View ดังภาพที่ 10.23 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ปริมาณเนื้อทุเรียน	Durian	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด
ปริมาณน้ำตาล	Sugar	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด
ค่าสี L^*	L	L^* Value	2	ไม่ต้องกำหนด

หมายเหตุ กรณีนี้กิจจจำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



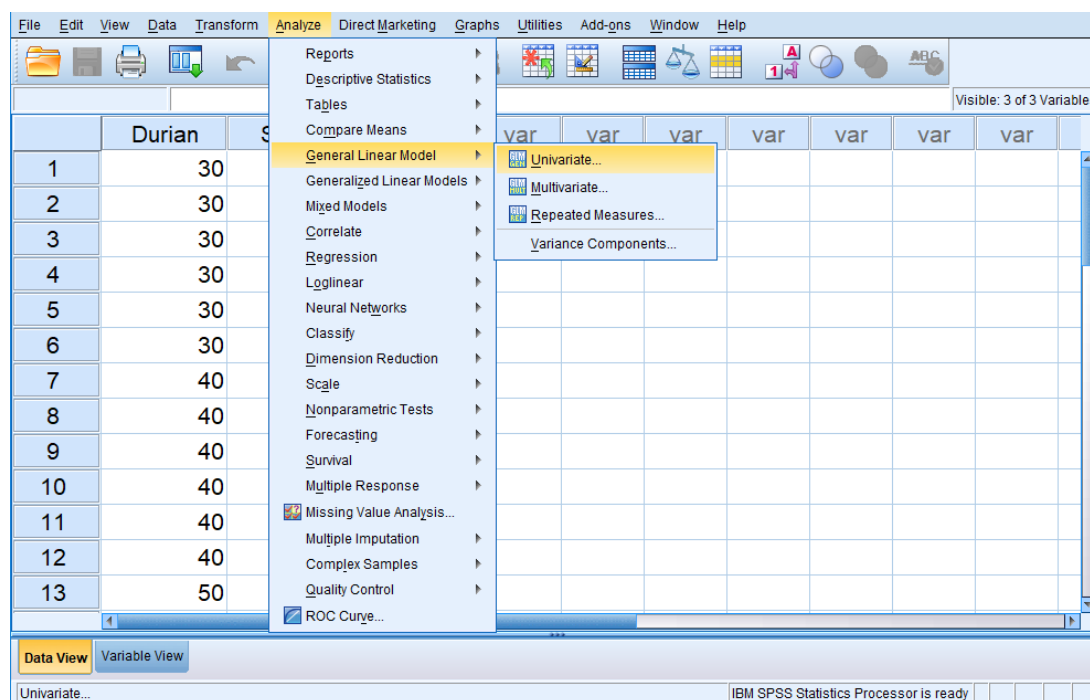
ภาพที่ 10.23 กำหนดรายละเอียดตัวแปร Durian, Sugar และ L ในหน้าต่าง Variable View

(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 10.24 โดยป้อนข้อมูลค่าสี L* ของเค้กทุเรียน แต่ละกลุ่มเรียงตามระดับปริมาณเนื้อทุเรียนและปริมาณน้ำตาล

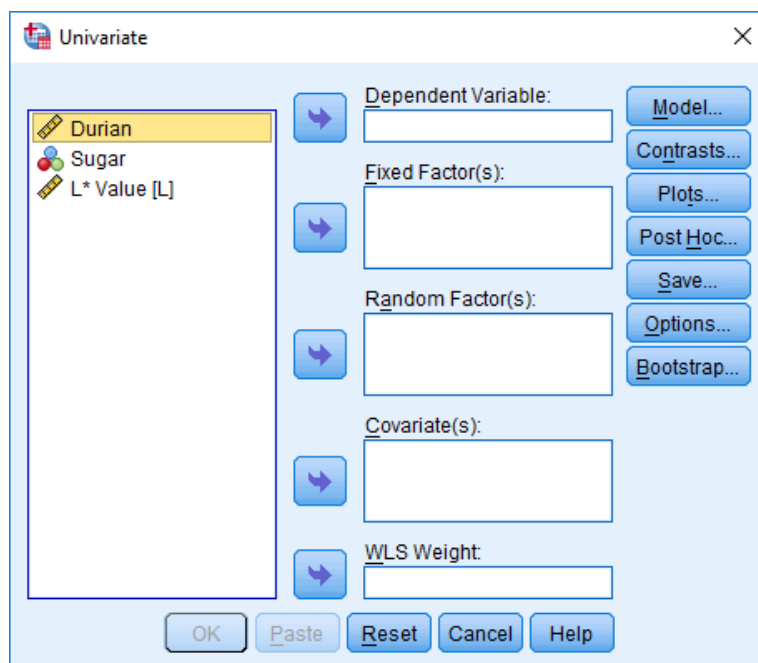
	Durian	Sugar	L	var	var	var	var	var	var	var
1	30	100	87.77							
2	30	100	87.17							
3	30	100	87.61							
4	30	150	79.83							
5	30	150	79.13							
6	30	150	80.07							
7	40	100	85.87							
8	40	100	85.67							
9	40	100	86.45							
10	40	150	75.65							
11	40	150	75.45							
12	40	150	75.81							
13	50	100	80.23							

ภาพที่ 10.24 ป้อนข้อมูลตัวแปร Durian, Sugar และ L ลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 10.25 จะปรากฏหน้าจอ ดังภาพที่ 10.26 จะเห็นว่าในกล่องซ้ายมือจะแสดงชื่อของตัวแปรแต่ละตัวที่กำหนดไว้

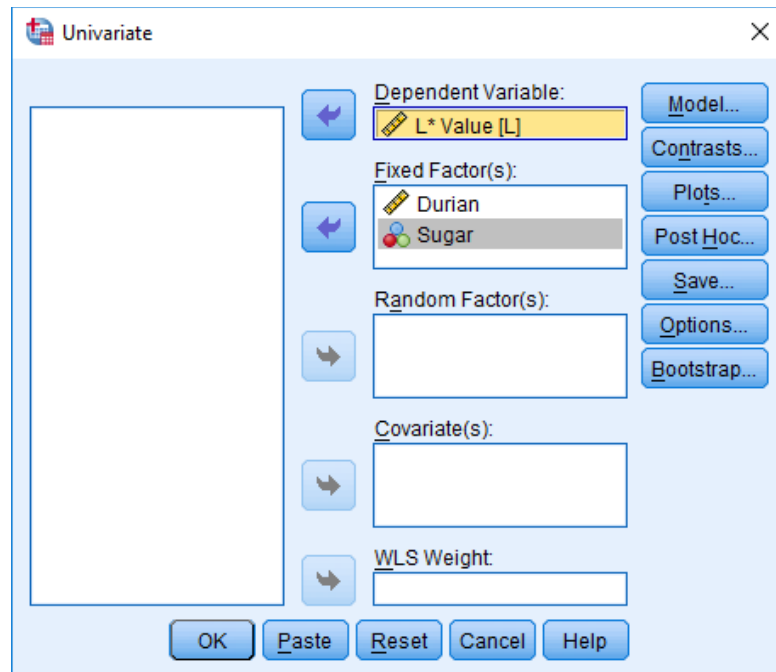


ภาพที่ 10.25 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



ภาพที่ 10.26 หน้าจอคำสั่ง Univariate

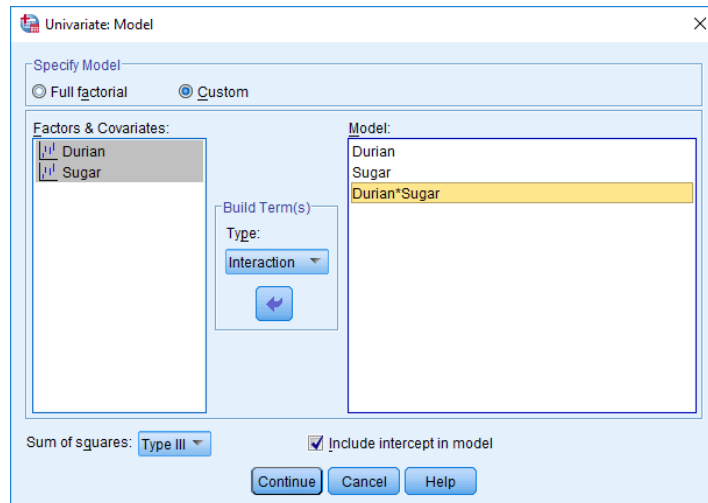
(4) เลือก ตัวแปร L*Value[L] (ตัวแปรตาม) ใส่ในช่อง **Dependent Variable** และ ตัวแปร Durian และ ตัวแปร Sugar (ตัวแปรต้น) ใส่ในช่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.27



ภาพที่ 10.27 เลือกตัวแปร L*Value[L] ใส่ในช่อง Dependent Variable และ ตัวแปร Durian และ ตัวแปร Sugar ใส่ในช่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Model...** จะได้หน้าจอตั้งภาพที่ 10.28 กำหนดค่าต่างๆ ดังนี้

- กำหนดรูปแบบการวิเคราะห์ (Model) ที่ตำแหน่ง Specify Model ให้คลิกเลือก **Custom** ทั้งนี้รูปแบบการวิเคราะห์ (Model) มี 2 รูปแบบ Full factorial คือ รูปแบบเต็มองค์ประกอบ และ Custom คือ การกำหนดรูปแบบขึ้นเอง
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลของปัจจัยหลัก (Main Effects) โดยกด Build Term(s) เลือก **Main Effects** แล้วเลือก ตัวแปร Durian และ ตัวแปร Sugar ในกล่องซ้ายมือใส่ในช่อง Model ที่อยู่ในกล่องขวามือ
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลร่วมระหว่างปัจจัยทั้ง 2 ปัจจัย (Interaction Effects) โดยกด Build Term(s) เลือก **Interaction** แล้วกดคีย์บอร์ด **Shift ↑** ค้างไว้ กดเลือกตัวแปร Durian และ ตัวแปร Sugar (กดเลือกแบบจับคู่) ในกล่องซ้ายมือใส่ในช่อง Model จะได้ตัวแปร **Durian*Sugar** ในกล่องขวามือ
- เมื่อกำหนดค่าเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



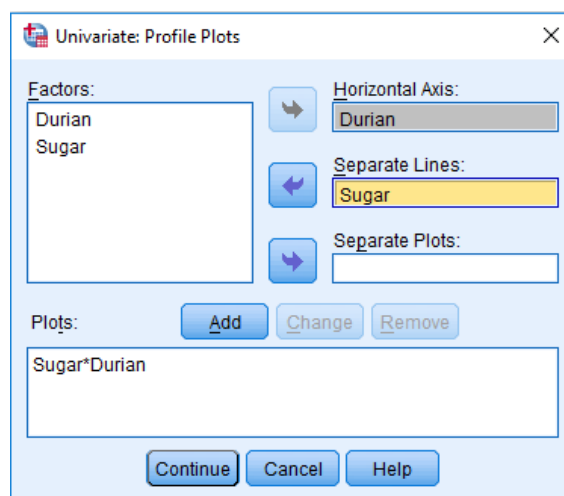
ภาพที่ 10.28 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลร่วมของทั้งสองปัจจัย จะได้หน้าจอตั้งภาพที่ 10.29 ในตัวอย่างนี้จะสร้าง 2 กราฟ โดยสลับตัวแปรที่ใช้เป็นแกนนอนและที่ใช้แสดงรูปแบบเส้นกราฟ

- กราฟที่ 1 กดเลือกตัวแปร Sugar ในช่อง Horizontal Axis (แกนนอน) และตัวแปร Durian ในช่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Sugar*Durian** ในช่องด้านล่าง

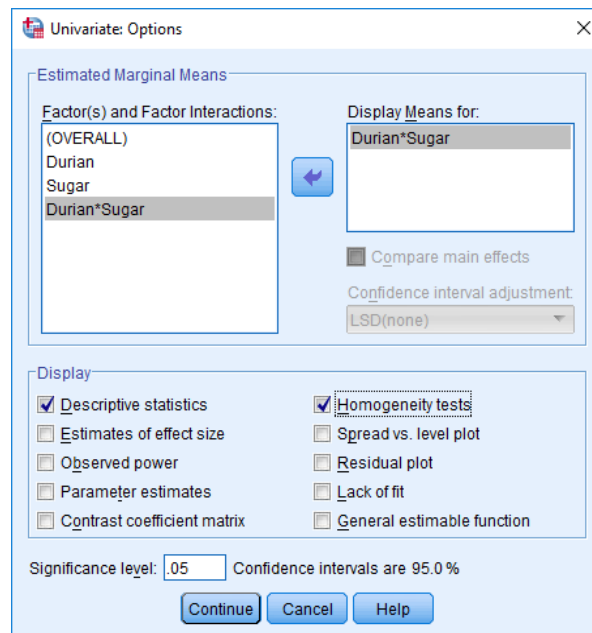
- กราฟที่ 2 กดเลือกตัวแปร Durian ในช่อง Horizontal Axis (แกนนอน) และตัวแปร Sugar ในช่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Durian*Sugar** ในช่องด้านล่าง

- เมื่อกำหนดตัวแปรเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 10.29 หน้าจอคำสั่ง Univariate : Profile Plots

(7) เลือกคำสั่ง **Options...** เพื่อให้แสดงค่าสถิติเบื้องต้นของแต่ละกลุ่ม จะได้หน้าจอตั้งภาพที่ 10.30 กดเลือก **ตัวแปร Durian*Sugar** ใส่ในช่อง Display Means for แล้วกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) เมื่อเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 10.31



ภาพที่ 10.30 หน้าจอคำสั่ง Univariate : Options

Univariate Analysis of Variance

Between-Subjects Factors

	N
Durian 30	6
40	6
50	6
Sugar 100	9
150	9

Descriptive Statistics

Dependent Variable: L*value

Durian	Sugar	Mean	Std. Deviation	N
30	100	87.5167	.31070	3
	150	79.6767	.48840	3
	Total	83.5967	4.30972	6
40	100	85.9967	.40513	3
	150	75.6367	.18037	3
	Total	80.8167	5.68133	6
50	100	80.2800	.17059	3
	150	70.2333	.22008	3
	Total	75.2567	5.50560	6
Total	100	84.5978	3.31549	9
	150	75.1822	4.11298	9
Total	Total	79.8900	6.04984	18

Levene's Test of Equality of Error Variances^a

Dependent Variable: L*value

F	df1	df2	Sig.
1.620	5	12	.228

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + Durian + Sugar + Durian * Sugar

Tests of Between-Subjects Effects

Dependent Variable: L*value

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	620.991 ^a	5	124.198	1223.093	.000
Intercept	114883.418	1	114883.418	1131360.937	.000
Durian	216.395	2	108.198	1065.520	.000
Sugar	398.937	1	398.937	3928.694	.000
Durian * Sugar	5.659	2	2.829	27.865	.000
Error	1.219	12	.102		
Total	115505.628	18			
Corrected Total	622.210	17			

a. R Squared = .998 (Adjusted R Squared = .997)

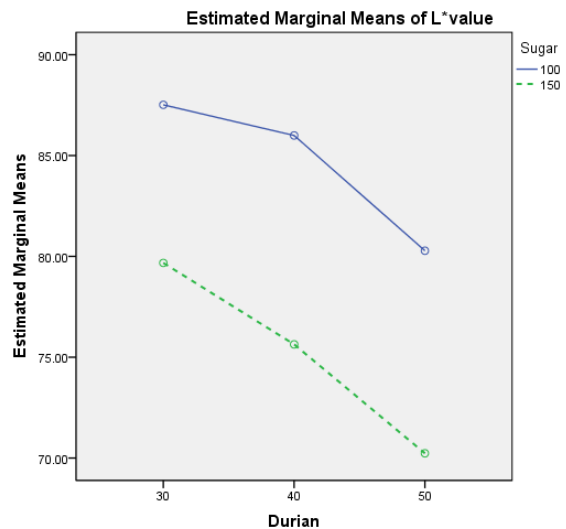
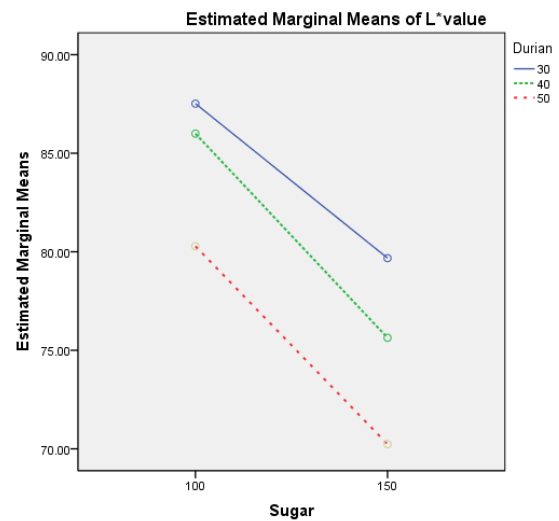
Estimated Marginal Means

Durian * Sugar

Dependent Variable: L*value

Durian	Sugar	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
30	100	87.517	.184	87.116	87.918
	150	79.677	.184	79.276	80.078
40	100	85.997	.184	85.596	86.398
	150	75.637	.184	75.236	76.038
50	100	80.280	.184	79.879	80.681
	150	70.233	.184	69.832	70.634

Profile Plots



ภาพที่ 10.31 ผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 3x2 factorial in CRD

10.5.4 การสรุปผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 3x2 Factorial in CRD

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่ได้กำหนดไว้ ดังนี้

สมมติฐานข้อที่ 1 ทดสอบอิทธิพลของระดับปริมาณเนื้อทุเรียน (Main effect A)

H_0 : เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30, 40 และ 50 กรัม มีค่าเฉลี่ยสี L^* ไม่แตกต่างกัน

H_1 : มีปริมาณเนื้อทุเรียน 2 ระดับที่เค้กทุเรียนมีค่าเฉลี่ยสี L^* แตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 10.32 ค่า Sig. ของตัวแปร Durian มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ มีปริมาณเนื้อทุเรียน 2 ระดับที่เค้กทุเรียนมีค่าเฉลี่ยสี L^* แตกต่างกัน ที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 2 ทดสอบอิทธิพลของระดับปริมาณน้ำตาล (Main effect B)

H_0 : เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยสี L^* ไม่แตกต่างกัน

H_1 : เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยสี L^* แตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 10.32 ค่า Sig. ของตัวแปร Sugar มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยสี L^* แตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 3 ทดสอบอิทธิพลร่วมระหว่างระดับปริมาณเนื้อทุเรียนและปริมาณน้ำตาล

(Interaction effects AxB)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยสี L^* ของเค้ก

H_1 : มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยสี L^* ของเค้ก

จากผลลัพธ์ที่ได้ในภาพที่ 10.32 ค่า Sig. ของตัวแปร Durian*Sugar มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยสี L^* ของเค้กที่ระดับนัยสำคัญ 0.05 และเมื่อพิจารณากราฟ Profile Plots พบว่า ในทุกระดับของปริมาณน้ำตาลเมื่อเพิ่มปริมาณเนื้อทุเรียนมีผลทำให้ค่าสี L^* ลดลง และในทุกระดับของปริมาณเนื้อทุเรียนเมื่อเพิ่มระดับน้ำตาลพบว่าค่าสี L^* มีค่าลดลง

10.5.5 การนำเสนอผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 3×2 Factorial in CRD

การนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 3×2 Factorial in CRD จะเป็นการนำเสนอผลลัพธ์เช่นเดียวกับแบบการวางแผนการทดลองแบบ 2² Factorial in CRD (ข้อ 10.4.4.6) ในที่นี้จะนำเสนอในรูปแบบตาราง 2 รูปแบบดังตารางที่ 10.9 และ 10.10 ซึ่งนักวิจัยจะเลือกนำเสนอแบบใดนั้นขึ้นอยู่กับวัตถุประสงค์ที่ต้องการนำเสนอข้อมูล

ตารางที่ 10.9 ค่า P-value แสดงอิทธิพลของปริมาณเนื้อทุเรียน อิทธิพลของปริมาณน้ำตาล และอิทธิพลร่วมระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยสี L* ของเค้กทุเรียน

อิทธิพลของปัจจัย	ค่า P-value
ปริมาณเนื้อทุเรียน	0.000*
ปริมาณน้ำตาล	0.000*
ปริมาณเนื้อทุเรียน × ปริมาณน้ำตาล	0.000*

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.10 ค่าเฉลี่ยสี L* ของเค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่างกัน

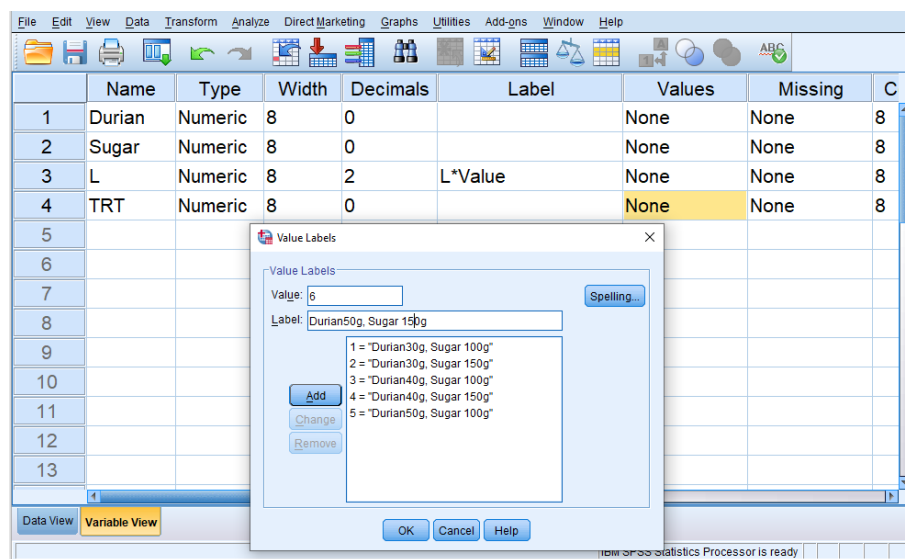
ปริมาณเนื้อทุเรียน (กรัม)	ปริมาณน้ำตาล (กรัม)	ค่าสี L*
30	100	87.52 ± 0.31a
30	150	79.68 ± 0.49d
40	100	86.00 ± 0.41b
40	150	75.64 ± 0.18e
50	100	80.28 ± 0.17c
50	150	70.23 ± 0.22f
อิทธิพลของปัจจัย		ค่า P-value
ปริมาณเนื้อทุเรียน		0.000*
ปริมาณน้ำตาล		0.000*
ปริมาณเนื้อทุเรียน × ปริมาณน้ำตาล		0.000*

^{a-f} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.10 ได้มาจากการกำหนดตัวแปร var ใหม่ชื่อ **TRT** ซึ่งเป็นตัวแปรที่แสดงถึงสิ่งทดลองที่เกิดจากปัจจัยร่วมระหว่างปริมาณเนื้อทุเรียน (3 ระดับ) และ ปริมาณน้ำตาล (2 ระดับ) หรือเรียกว่า **“Treatment Combination”** ในตัวอย่างนี้จะได้จำนวนสิ่งทดลอง (TRT) เท่ากับ 6 สิ่งทดลอง (คำนวณจากผลคูณของระดับแต่ละปัจจัย) นำข้อมูลตัวแปร TRT และตัวแปร L ไปวิเคราะห์ด้วยคำสั่ง **Post Hoc...** เพื่อเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ ทำเช่นเดียวกับการวิเคราะห์ความแปรปรวนแบบทางเดียว (One-Way ANOVA) ตามขั้นตอนดังนี้

(1) กำหนดรายละเอียดและค่า Values ของตัวแปร TRT ในหน้าต่าง Variable View ดังภาพที่ 10.32 และป้อนข้อมูล TRT ในหน้าต่าง Data View ดังภาพที่ 10.33

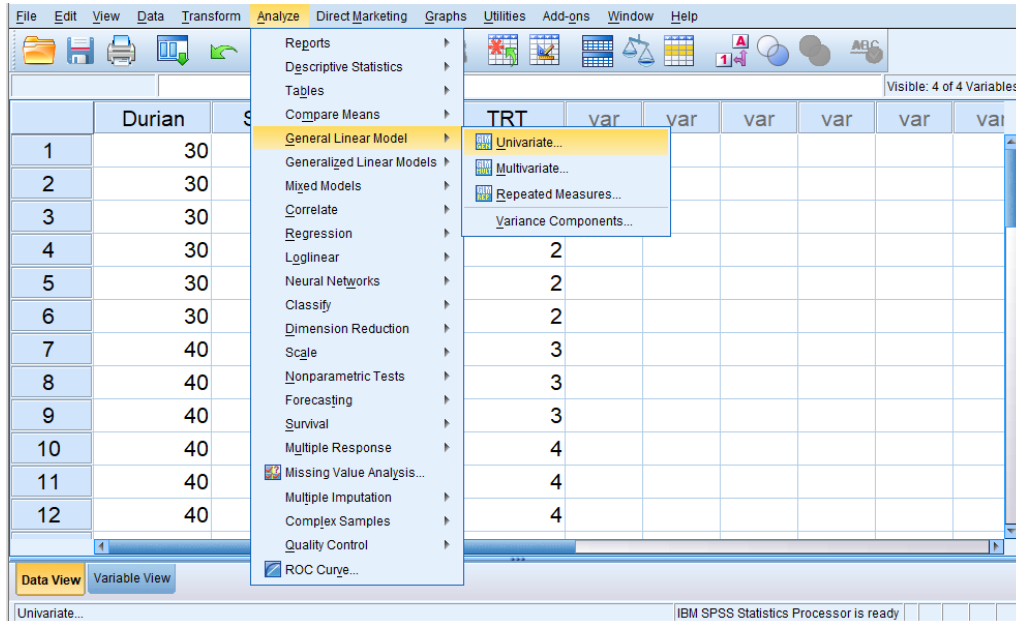


ภาพที่ 10.32 การกำหนดรายละเอียดตัวแปร TRT ในหน้าต่าง Variable View

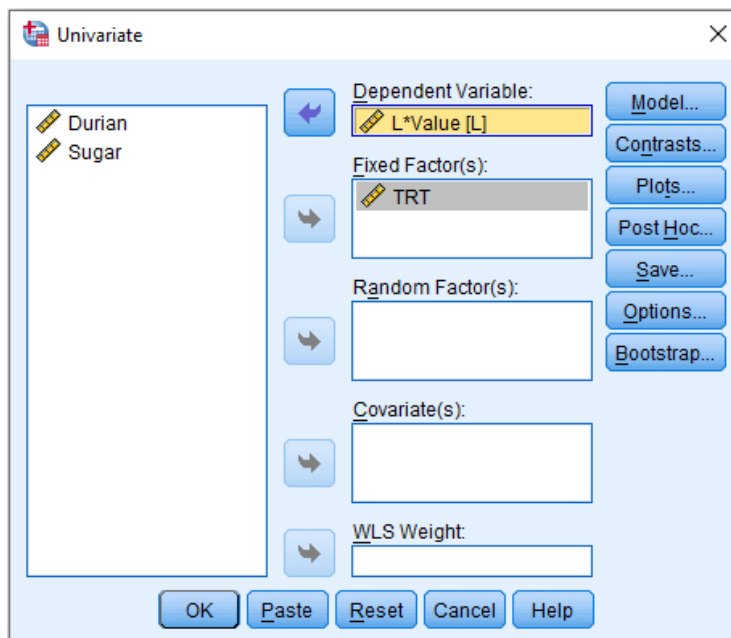
	Durian	Sugar	L	TRT	var	var	var	var	var	var
1	30	100	87.77	1						
2	30	100	87.17	1						
3	30	100	87.61	1						
4	30	150	79.83	2						
5	30	150	79.13	2						
6	30	150	80.07	2						
7	40	100	85.87	3						
8	40	100	85.67	3						
9	40	100	86.45	3						
10	40	150	75.65	4						
11	40	150	75.45	4						
12	40	150	75.81	4						

ภาพที่ 10.33 การป้อนข้อมูล TRT ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 10.34 และเลือกตัวแปร **L*Value [L]** (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร **TRT** (ตัวแปรต้น) ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.35

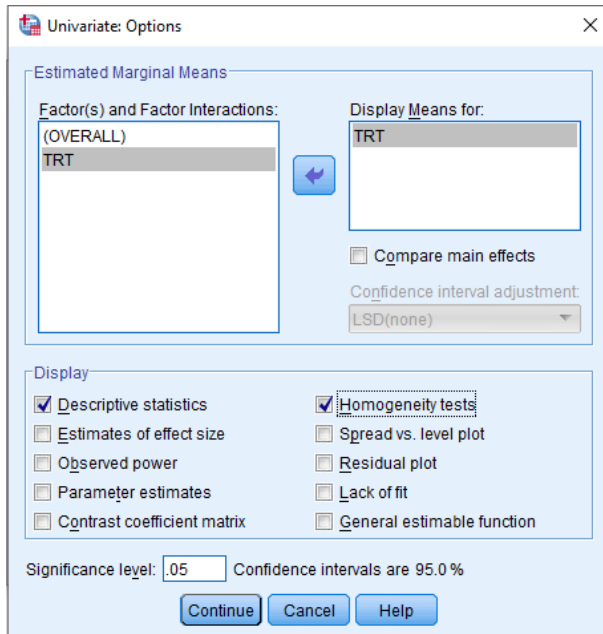


ภาพที่ 10.34 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



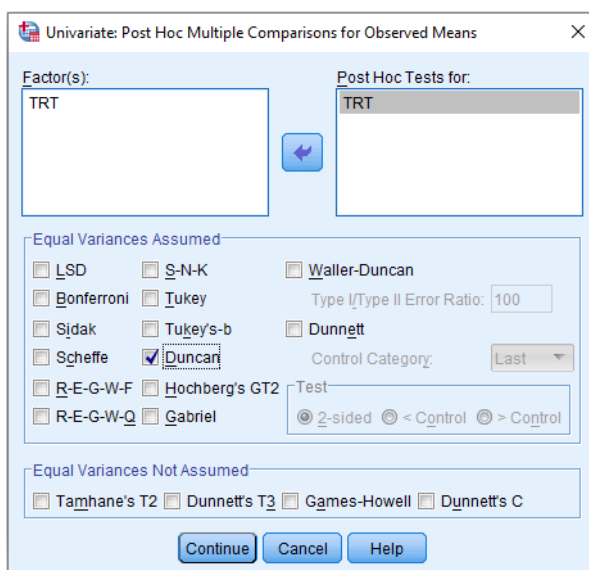
ภาพที่ 10.35 การกำหนดตัวแปรในเมนูคำสั่ง Univariate

(3) เลือกคำสั่ง **Options...** จะได้หน้าจอ ดังภาพที่ 10.36 กดเลือก **ตัวแปร TRT** ใส่ในกล่อง Display Means for ที่อยู่ด้านขวามือ และกดเลือก Homogeneity tests (เพื่อทดสอบความแปรปรวนของข้อมูลแต่ละสิ่งทดลอง) และกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 10.36 การกำหนดตัวแปรในเมนูคำสั่ง Univariate : Options

(5) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอ ดังภาพที่ 10.37 ให้กดเลือก **ตัวแปร TRT** ใส่ในกล่อง Post Hoc Tests for: และกดเลือก **วิธี Duncan** ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 10.38



ภาพที่ 10.37 การกำหนดตัวแปรและวิธีการทดสอบในเมนูคำสั่ง Post Hoc...

Univariate Analysis of Variance

Between-Subjects Factors			
TRT		Value Label	N
1	1	Durian30g, Sugar 100g	3
	2	Durian30g, Sugar 150g	3
	3	Durian40g, Sugar 100g	3
	4	Durian40g, Sugar 150g	3
	5	Durian50g, Sugar 100g	3
	6	Durian50g, Sugar 150g	3

Tests of Between-Subjects Effects					
Dependent Variable: L*Value					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	620.991 ^a	5	124.198	1223.093	.000
Intercept	114883.418	1	114883.418	1131360.937	.000
TRT	620.991	5	124.198	1223.093	.000
Error	1.219	12	.102		
Total	115505.628	18			
Corrected Total	622.210	17			

a. R Squared = .998 (Adjusted R Squared = .997)

Estimated Marginal Means

Descriptive Statistics			
Dependent Variable: L*Value			
TRT	Mean	Std. Deviation	N
Durian30g, Sugar 100g	87.5167	.31070	3
Durian30g, Sugar 150g	79.6767	.48840	3
Durian40g, Sugar 100g	85.9967	.40513	3
Durian40g, Sugar 150g	75.6367	.18037	3
Durian50g, Sugar 100g	80.2800	.17059	3
Durian50g, Sugar 150g	70.2333	.22008	3
Total	79.8900	6.04984	18

TRT				
Dependent Variable: L*Value				
TRT	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Durian30g, Sugar 100g	87.517	.184	87.116	87.918
Durian30g, Sugar 150g	79.677	.184	79.276	80.078
Durian40g, Sugar 100g	85.997	.184	85.596	86.398
Durian40g, Sugar 150g	75.637	.184	75.236	76.038
Durian50g, Sugar 100g	80.280	.184	79.879	80.681
Durian50g, Sugar 150g	70.233	.184	69.832	70.634

Levene's Test of Equality of Error Variances ^a			
Dependent Variable: L*Value			
F	df1	df2	Sig.
1.620	5	12	.228

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept + TRT

Post Hoc Tests

TRT

Homogeneous Subsets

L*Value							
Duncan ^{a,b}							
TRT	N	Subset					
		1	2	3	4	5	6
Durian50g, Sugar 150g	3	70.2333	75.6367	79.6767	80.2800	85.9967	87.5167
Durian40g, Sugar 150g	3						
Durian30g, Sugar 150g	3						
Durian50g, Sugar 100g	3						
Durian40g, Sugar 100g	3						
Durian30g, Sugar 100g	3						
Sig.		1.000	1.000	1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .102.

a. Uses Harmonic Mean Sample Size = 3.000.
b. Alpha = .05.

ภาพที่ 10.38 ผลลัพธ์การวิเคราะห์ความแปรปรวนของตัวแปร TRT และ ตัวแปร L

10.6 การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย

ในบางงานวิจัยเป็นไปได้ที่จะทำการทดลองแบบแฟกทอเรียลโดยวางแผนการทดลองแบบสุ่มสมบูรณ์ (CRD) ตัวอย่างเช่น งานวิจัยที่มีปัจจัยรบกวนหรือปัจจัยที่เกี่ยวข้อง (Nuisance factor) เช่น สถานที่ทดสอบต่างกัน เวลาต่างกัน แหล่งวัตถุดิบต่างกัน เพศต่างกัน จึงจำเป็นต้องทำการทดลองโดยการแบ่งเป็นบล็อก โดยเฉพาะในงานวิจัยและพัฒนาผลิตภัณฑ์ที่มีผู้ชมเข้ามาเกี่ยวข้องในการประเมินคุณภาพทางประสาทสัมผัส เรามักใช้การวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (Randomized Complete Block Design หรือ RCBD) โดยกำหนดให้ผู้ชม คือ บล็อก (อ่านรายละเอียดการวางแผนการทดลองแบบ RCBD เพิ่มเติมได้ในบทที่ 9)

การวางแผนการทดลองแบบ Factorial in RCBD ชนิดที่มี 2 ปัจจัย จัดเป็นการทดลองที่ง่ายที่สุด เป็นการจัดสิ่งทดลองที่มี 2 ปัจจัย ได้แก่ ปัจจัย A มี a ระดับ และปัจจัย B มี b ระดับ เมื่อนำระดับมาจัดกลุ่ม สิ่งทดลองจะได้จำนวนสิ่งทดลองที่เป็นไปได้ เท่ากับ ab สิ่งทดลอง (ระดับของปัจจัย A $\times$ ระดับของปัจจัย B) และถ้ามีการบล็อกในแต่ละสิ่งทดลอง n ครั้ง จะได้จำนวนหน่วยทดลอง (Experimental unit) หรือค่าสังเกตทั้งหมดเท่ากับ abn หน่วย กรณีที่ปัจจัยเป็นอิทธิพลที่จะเขียนตัวแบบ (Model) ของแผนการทดลองแบบ Factorial in RCBD ชนิดที่มี 2 ปัจจัย ได้ดังนี้

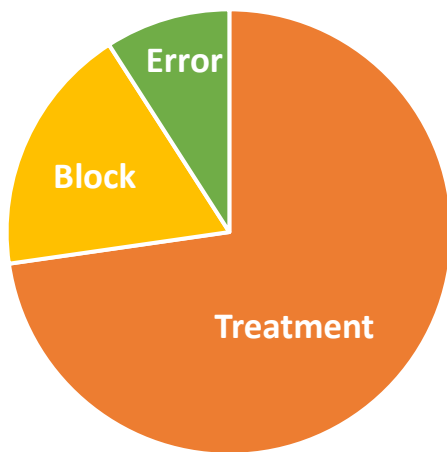
$$y_{ijk} = \mu + \tau_i + \beta_j + (\tau\beta)_{ij} + \delta_k + \varepsilon_{ijk} \quad ; \quad \begin{array}{l} i = 1, 2, \dots, a \\ j = 1, 2, \dots, b \\ k = 1, 2, \dots, n \end{array}$$

โดยที่ y_{ijk} คือ ค่าสังเกตที่บล็อก k ของปัจจัย A ที่ระดับ i และปัจจัย B ที่ระดับ j
 μ คือ ค่าเฉลี่ยของข้อมูลทั้งหมด
 τ_i คือ อิทธิพล (Effect) จากปัจจัย A ที่ระดับ i
 β_j คือ อิทธิพล (Effect) จากปัจจัย B ที่ระดับ j
 $(\tau\beta)_{ij}$ คือ อิทธิพล (Effect) ของปฏิสัมพันธ์ระหว่าง τ_i และ β_j
 δ_k คือ อิทธิพล (Effect) จากบล็อก ที่ระดับ k
 ε_{ijk} คือ ส่วนประกอบความคลาดเคลื่อนแบบสุ่ม (Experimental error หรือ Random error)

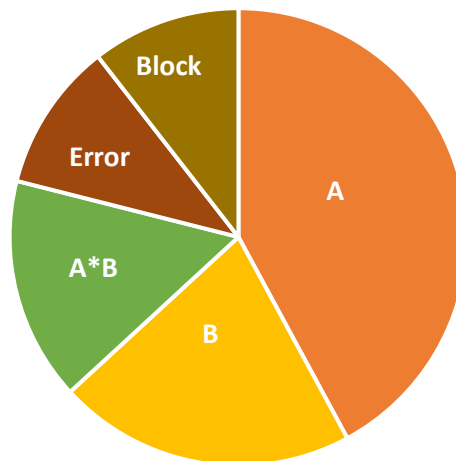
10.6.1 ความแปรปรวนของการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย

โดยทั่วไปแล้วการวางแผนการทดลองแบบสุ่มในบล็อกสมบูรณ์ (RCBD) เป็นการวางแผนการทดลองที่มี 1 ปัจจัย ดังนั้นจึงมีแหล่งของความแปรปรวนที่เกิดขึ้น 3 แหล่ง แบ่งเป็น ความแปรปรวนที่อธิบายได้ 2 แหล่งจากสิ่งทดลอง (Treatment) จากบล็อก (Block) และ ความแปรปรวนที่อธิบายไม่ได้ 1 แหล่งจากความคลาดเคลื่อนของการทดลอง (Experimental error) ดังภาพที่ 10.39(a) และเงื่อนไขของการวางแผนการทดลองแบบ RCBD คือ ต้องไม่มีปฏิสัมพันธ์ (interaction) ระหว่างบล็อกและสิ่งทดลอง แต่ในกรณีที่วางแผนการทดลองแบบ Factorial in RCBD และมีปัจจัยที่ต้องการศึกษา 2 ปัจจัย (ปัจจัย A และปัจจัย B) นั้นหมายความว่า สิ่งทดลองเกิดจากการจัดกลุ่มของระดับสองปัจจัยร่วมกัน ดังนั้น ความแปรปรวนของสิ่งทดลองจึงประกอบด้วย ความแปรปรวนที่เกิดจากปัจจัย A, ความแปรปรวนที่เกิดจากปัจจัย B และความแปรปรวนที่เกิดจากปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B

กล่าวโดยสรุปคือ แหล่งของความแปรปรวนในการวางแผนการทดลองแบบ Factorial in RCBD ชนิดที่มี 2 ปัจจัย จะประกอบด้วย 5 แหล่ง แบ่งเป็น ความแปรปรวนที่อธิบายได้ 4 แหล่ง จาก ปัจจัย A, ปัจจัย B, ปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B และบล็อก และ ความแปรปรวนที่อธิบายไม่ได้ 1 แหล่ง จากความคลาดเคลื่อนของการทดลอง ดังภาพที่ 10.39(b) การวิเคราะห์ความแปรปรวนจะเป็นแบบสองทาง (Two-way ANOVA) ซึ่งสามารถแสดงผลตารางการวิเคราะห์ความแปรปรวนได้ดังตารางที่ 10.11



(a) การวางแผนการทดลองแบบ RCBD



(b) การวางแผนการทดลองแบบ

Factorial in RCBD ชนิดที่มี 2 ปัจจัย

ภาพที่ 10.39 เปรียบเทียบการแบ่งสัดส่วนความแปรปรวนของ (a) การวางแผนการทดลองแบบ RCBD และ (b) การวางแผนการทดลองแบบ Factorial in RCBD ชนิดที่มี 2 ปัจจัย

ตารางที่ 10.11 การวิเคราะห์ความแปรปรวนสำหรับการวางแผนการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย โดยมีตัวแบบอิทธิพลแบบคงที่ (Fixed effect model)

แหล่งความแปรปรวน (Source of Variation)	องศาความเป็นอิสระ (Degree of Freedom)	ผลบวกกำลังสอง (Sum of Squares)	ค่ากำลังสองเฉลี่ย (Mean Squares)	F
บล็อก	$n - 1$	SSBlock	$MSBlock = \frac{SSBlock}{n - 1}$	$F = \frac{MSBlock}{MSE}$
ปัจจัย A	$a - 1$	SSA	$MSA = \frac{SSA}{a - 1}$	$F = \frac{MSA}{MSE}$
ปัจจัย B	$b - 1$	SSB	$MSB = \frac{SSB}{b - 1}$	$F = \frac{MSB}{MSE}$
ปฏิสัมพันธ์ AB	$(a - 1)(b - 1)$	SSAB	$MSAB = \frac{SSAB}{(a - 1)(b - 1)}$	$F = \frac{MSAB}{MSE}$
ความคลาดเคลื่อน	$(ab - 1)(n - 1)$	SSE	$MSE = \frac{SSE}{(ab - 1)(n - 1)}$	
ยอดรวม (Total)	$abn - 1$	SST		

a = จำนวนสิ่งทดลองหรือระดับของปัจจัย A

b = จำนวนสิ่งทดลองหรือระดับของปัจจัย B

n = จำนวนหรือระดับของ Block

โดยทั่วไปแล้ว ขั้นตอนแรกของการวิเคราะห์ความแปรปรวนจะทำการทดสอบอิทธิพลของบล็อกและปฏิสัมพันธ์ระหว่างปัจจัยก่อน หลังจากนั้นจึงทำการทดสอบอิทธิพลหลัก ข้อสังเกต คือ

(1) ถ้าอิทธิพลของบล็อกมีนัยสำคัญทางสถิติ แสดงว่าการเลือกวางแผนการทดลองแบบ RCBD ถูกต้องแล้ว แต่ถ้าอิทธิพลของบล็อกไม่มีนัยสำคัญทางสถิติควรพิจารณาให้รอบคอบว่าในการทดลองครั้งต่อไปสมควรเปลี่ยนไปใช้การวางแผนการทดลองแบบ CRD แทนหรือไม่ (อ่านรายละเอียดเพิ่มเติมในบทที่ 8 และ 9)

(2) ถ้าปฏิสัมพันธ์ระหว่างปัจจัยไม่มีนัยสำคัญทางสถิติ ให้ทำการแปลความหมายของการทดสอบอิทธิพลหลักได้โดยตรง แต่ถ้าปฏิสัมพันธ์มีนัยสำคัญทางสถิติ อิทธิพลหลักของปัจจัยต่างๆ จะเกี่ยวข้องกับปฏิสัมพันธ์ ซึ่งอาจจะไม่มีประโยชน์ในการแปลความหมายของการทดสอบอิทธิพลหลักในทางปฏิบัติมากนัก โดยทั่วไปความรู้เกี่ยวกับปฏิสัมพันธ์มีความสำคัญมากกว่าความรู้เกี่ยวกับอิทธิพลหลัก

10.6.2 การเขียนสมมติฐานทางสถิติสำหรับการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย

ในการทดสอบสมมติฐานการวางแผนการทดลองแบบ Factorial in RCBD ชนิดที่มี 2 ปัจจัย คือ ปัจจัย A และ ปัจจัย B จะให้ความสนใจในการทดสอบอิทธิพลของปัจจัยหลัก ได้แก่ อิทธิพลของปัจจัย A และ อิทธิพลของปัจจัย B ที่มีต่อค่าสังเกต และอิทธิพลร่วมระหว่างระดับของปัจจัย A และปัจจัย B ที่มีต่อค่าสังเกต อย่างไรก็ตามสามารถทดสอบอิทธิพลของบล็อกได้เช่นกัน ดังนั้นการทดสอบสมมติฐานจะประกอบด้วย

สมมติฐานข้อ 1 การทดสอบอิทธิพลของปัจจัย A (Main effect A) ต่อตัวแปรตาม

H_0 : ปัจจัย A ไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : ปัจจัย A มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ H_0 : ค่าเฉลี่ยของตัวอย่างในทุกะดับของปัจจัย A ไม่แตกต่างกัน

H_1 : มีระดับของปัจจัย A อย่างน้อย 2 ระดับที่มีค่าเฉลี่ยแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

$H_0 : \tau_i = 0$

$H_1 : \tau_i \neq 0$ อย่างน้อย 1 ค่า

สมมติฐานข้อ 2 การทดสอบอิทธิพลของปัจจัย B (Main effect B) ต่อตัวแปรตาม

H_0 : ปัจจัย B ไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : ปัจจัย B มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ H_0 : ค่าเฉลี่ยของตัวอย่างในทุกะดับของปัจจัย B ไม่แตกต่างกัน

H_1 : มีระดับของปัจจัย B อย่างน้อย 2 ระดับที่มีค่าเฉลี่ยแตกต่างกัน

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

$H_0 : \beta_j = 0$

$H_1 : \beta_j \neq 0$ อย่างน้อย 1 ค่า

สมมติฐานข้อ 3 การทดสอบอิทธิพลร่วมระหว่างปัจจัย A และปัจจัย B (Interaction effects A×B) ต่อตัวแปรตาม

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B ต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : มีปฏิสัมพันธ์ระหว่างปัจจัย A และปัจจัย B ต่อค่าเฉลี่ยของตัวแปรตาม

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

$H_0 : (\tau\beta)_{ij} = 0$

$H_1 : (\tau\beta)_{ij} \neq 0$ อย่างน้อย 1 ค่า

สมมติฐานข้อ 4 การทดสอบอิทธิพลของบล็อก (Block) ต่อตัวแปรตาม

H_0 : บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

H_1 : บล็อกมีอิทธิพลต่อค่าเฉลี่ยของตัวแปรตาม

หรือ เขียนเชิงสัญลักษณ์ทางคณิตศาสตร์ได้ ดังนี้

H_0 : $\delta_k = 0$

H_1 : $\delta_k \neq 0$ อย่างน้อย 1 ค่า

10.6.3 การทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ (Multiple Comparison Test)

เมื่อวิเคราะห์ความแปรปรวนแล้วสรุปผลว่า ปัจจัยหลักมีอิทธิพลต่อตัวแปรตาม และ/หรือ มีอิทธิพลร่วมระหว่างปัจจัยหลักต่อตัวแปรตาม นักวิจัยส่วนใหญ่มักต้องการวิเคราะห์เปรียบเทียบความแตกต่างของค่าเฉลี่ยในแต่ละสิ่งทดลองที่เกิดจากการร่วมระดับของปัจจัย (Treatment combination) ในกรณีนี้ต้องทำการทดสอบความแตกต่างของค่าเฉลี่ยแบบจับคู่พหุคูณ ยกตัวอย่างเช่น การวางแผนการทดลองแบบ 2^2 Factorial in RCBD มีปัจจัยที่ศึกษา 2 ปัจจัย และแต่ละปัจจัยมี 2 ระดับ จะมีจำนวนสิ่งทดลอง เท่ากับ 4 สิ่งทดลองให้ทำการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ

10.6.4 คำสั่ง SPSS ในการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ Factorial in RCBD ชนิด 2 ปัจจัย

คำสั่งที่ใช้ในการวิเคราะห์ข้อมูลการทดลองแบบ Factorial in RCBD ชนิดที่มี 2 ปัจจัย คือ การวิเคราะห์ความแปรปรวนแบบสองทาง (Two-Way ANOVA) ด้วยคำสั่ง ดังนี้

Analyze ➤ General Linear Model ➤ Univariate...

10.6.4.1 ตัวอย่างข้อมูลการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

นักวิจัยต้องการศึกษาว่าอุณหภูมิและเวลาในการทอดมีผลต่อคุณภาพของมันฝรั่งทอดในด้านกายภาพ เคมี และประสาทสัมผัสหรือไม่ จึงได้ทำการศึกษาอุณหภูมิในการทอด 2 ระดับ (150 และ 160 °C) และเวลาในการทอด 2 ระดับ (3 และ 5 นาที) ทั้งนี้เพื่อให้มีจำนวนตัวอย่างในแต่ละสิ่งทดลองเพียงพอสำหรับการวิเคราะห์คุณภาพดังกล่าว ในการทดลองจะต้องทอดมันฝรั่งจำนวน 3 รอบ ซึ่งการใช้น้ำมันทอดซ้ำกันอาจจะมีผลต่อค่าคุณภาพของมันฝรั่งทอดได้ ดังนั้นนักวิจัยจึงได้วางแผนการทดลองแบบ 2^2 Factorial in

RCBD โดยกำหนดให้ปัจจัย A คือ อุณหภูมิในการทอด และปัจจัย B คือ เวลาในการทอด และบล็อก คือ รอบของการทอด ทั้งนี้ในการวิเคราะห์ค่าคุณภาพต่างๆ นักวิจัยได้ทำการสุ่มตัวอย่างสิ่งทดลองในแต่ละบล็อกเพื่อมาวิเคราะห์คุณภาพ

จากการวางแผนการทดลองดังกล่าว นักวิจัยได้ทำการสุ่มตัวอย่างมาวิเคราะห์ปริมาณไขมันในห้องปฏิบัติการ ได้ข้อมูลดังตารางที่ 10.12

ตารางที่ 10.12 ปริมาณไขมันของมันฝรั่งทอดที่ใช้อุณหภูมิและเวลาในการทอดแตกต่างกัน

สิ่งทดลองที่	อุณหภูมิในการทอด (°C)	เวลาในการทอด (นาที)	ปริมาณไขมัน (%)		
			ทอดรอบ 1	ทอดรอบ 2	ทอดรอบ 3
1	150	3	27.6	27.5	27.6
2	150	5	26.5	26.6	26.0
3	160	3	25.7	25.9	26.1
4	160	5	24.4	24.6	24.8

10.6.4.2 การเขียนสมมติฐาน

จากข้อมูลดังกล่าว สามารถสรุปรายละเอียดของข้อมูลได้ ดังนี้

- ปัจจัย A คือ อุณหภูมิในการทอด มี 2 ระดับ (150 และ 160 °C)
- ปัจจัย B คือ เวลาในการทอด มี 2 ระดับ (3 และ 5 นาที)
- บล็อก คือ รอบในการทอด มี 3 บล็อก (ทอด 3 รอบ)
- จำนวนสิ่งทดลอง (Treatment) เท่ากับ 4 สิ่งทดลอง (ผลคูณของระดับแต่ละปัจจัย)
- จำนวนหน่วยทดลอง (Experimental unit) ในแต่ละสิ่งทดลอง เท่ากับ 3 หน่วย
- จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 12 หน่วย
- ค่าสังเกต หรือค่าตอบสนอง (ตัวแปรตาม) คือ ปริมาณไขมัน

นักวิจัยสามารถเขียนสมมติฐานที่ใช้ทดสอบ ได้ดังนี้

(1) การทดสอบสมมติฐานอิทธิพลของอุณหภูมิในการทอด (Main effect A)

H_0 : อุณหภูมิในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

(2) การทดสอบสมมติฐานอิทธิพลของเวลาในการทอด (Main effect B)

H_0 : เวลาในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

(3) การทดสอบสมมติฐานอิทธิพลร่วมระหว่างอุณหภูมิในการทอดและเวลาในการทอด

(Interaction effects AxB)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

(4) การทดสอบสมมติฐานอิทธิพลของบล็อก (Block)

H_0 : บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

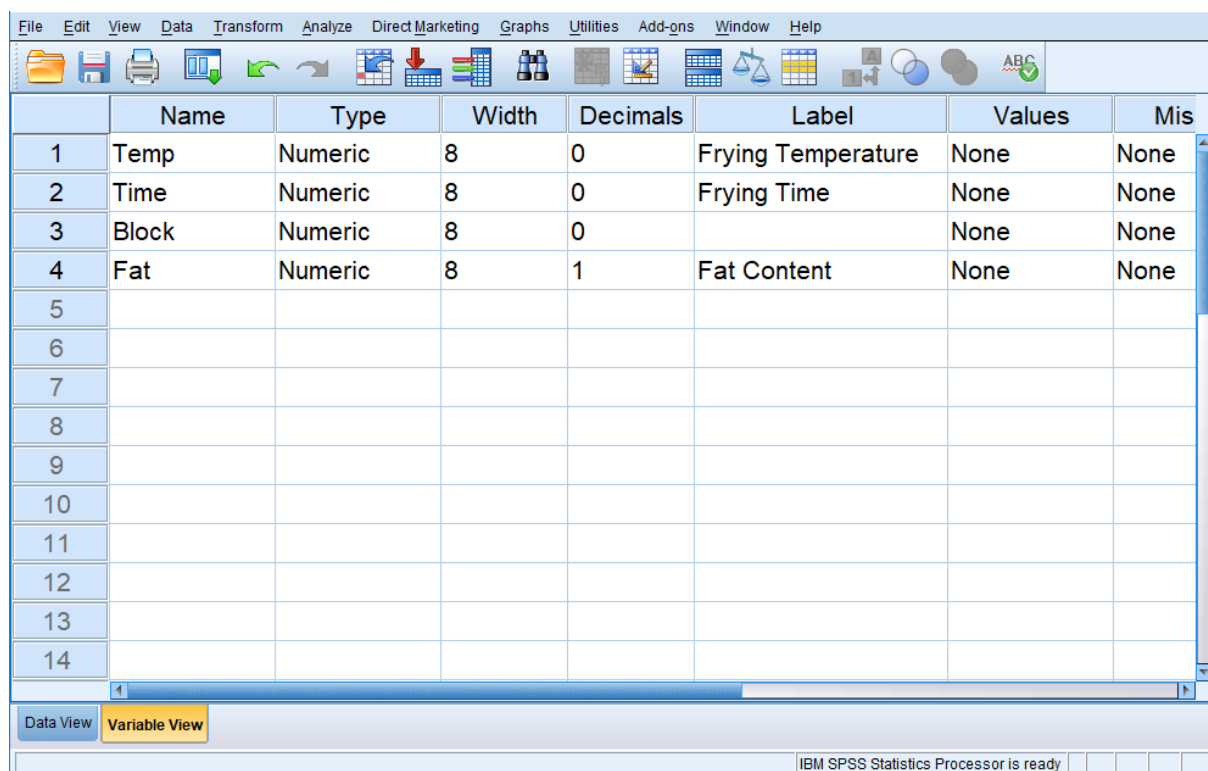
H_1 : บล็อกมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

10.6.4.3 การใช้คำสั่งในโปรแกรม SPSS วิเคราะห์ข้อมูลการทดลองแบบ 2^2 Factorial in RCBD

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลปริมาณไขมันของมันฝรั่งทอดแต่ละสิ่งทดลอง จะต้องใช้ตัวแปร var จำนวน 4 ตัวแปรสำหรับป้อนข้อมูล 4 ตัว ได้แก่ อุณหภูมิในการทอด, เวลาในการทอด, รอบในการทอด และปริมาณไขมัน ดังนั้นกำหนดชื่อตัวแปรเป็น **Temp**, **Time**, **Block** และ **Fat** ในหน้าต่าง Variable View ดังภาพที่ 10.40 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ระดับอุณหภูมิในการทอด	Temp	Frying Temperature	0	ไม่ต้องกำหนด
ระดับเวลาในการทอด	Time	Frying Time	0	ไม่ต้องกำหนด
รอบในการทอด	Block	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด
ปริมาณไขมัน	Fat	Fat Content	1	ไม่ต้องกำหนด

หมายเหตุ กรณีนี้กรณียกข้อยกเว้น (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้

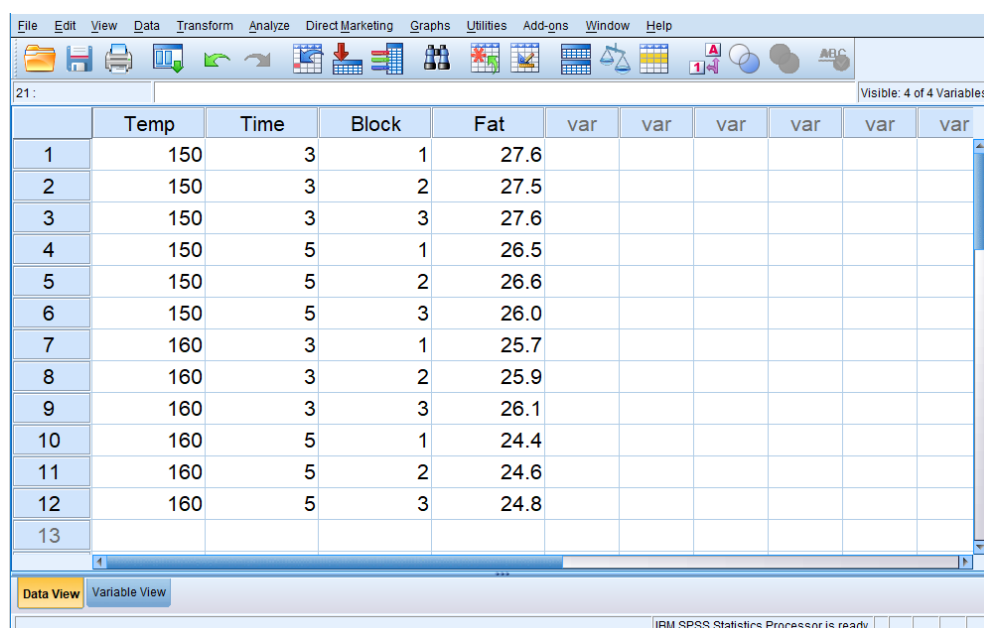


	Name	Type	Width	Decimals	Label	Values	Mis
1	Temp	Numeric	8	0	Frying Temperature	None	None
2	Time	Numeric	8	0	Frying Time	None	None
3	Block	Numeric	8	0		None	None
4	Fat	Numeric	8	1	Fat Content	None	None
5							
6							
7							
8							
9							
10							
11							
12							
13							
14							

IBM SPSS Statistics Processor is ready

ภาพที่ 10.40 กำหนดรายละเอียดตัวแปร Temp, Time, Block และ Fat ใน Variable View

(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 10.41 โดยป้อนข้อมูลปริมาณไขมันเรียงตามระดับอุณหภูมิและเวลาในการทอดและบล็อก

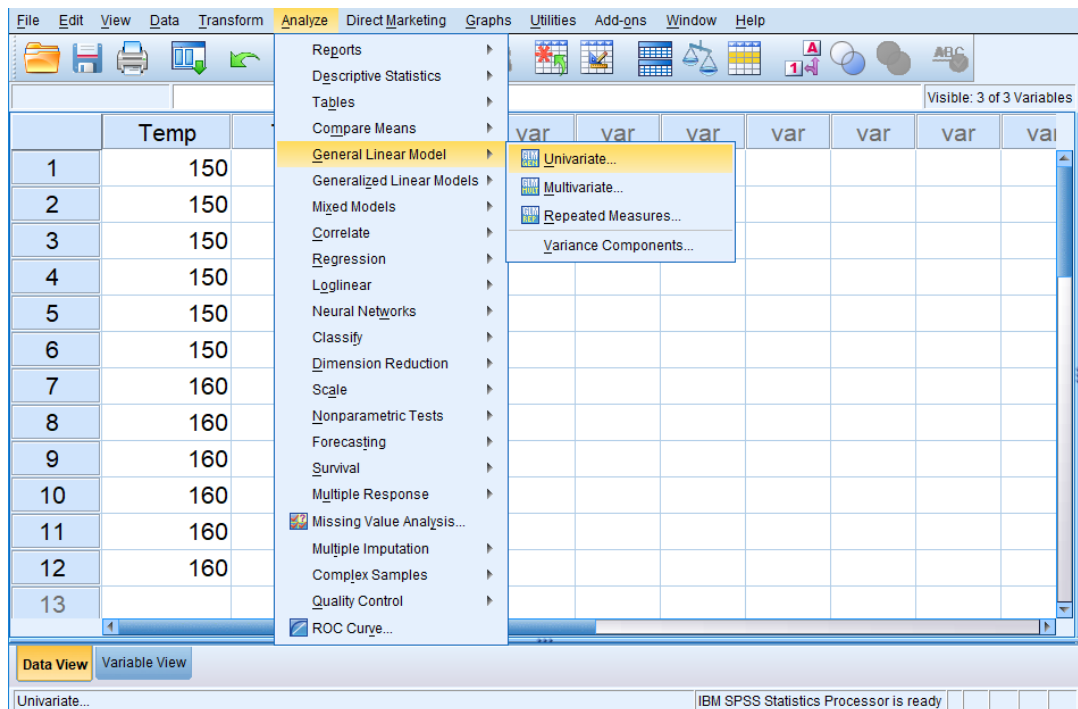


	Temp	Time	Block	Fat	var	var	var	var	var	var
1	150	3	1	27.6						
2	150	3	2	27.5						
3	150	3	3	27.6						
4	150	5	1	26.5						
5	150	5	2	26.6						
6	150	5	3	26.0						
7	160	3	1	25.7						
8	160	3	2	25.9						
9	160	3	3	26.1						
10	160	5	1	24.4						
11	160	5	2	24.6						
12	160	5	3	24.8						
13										

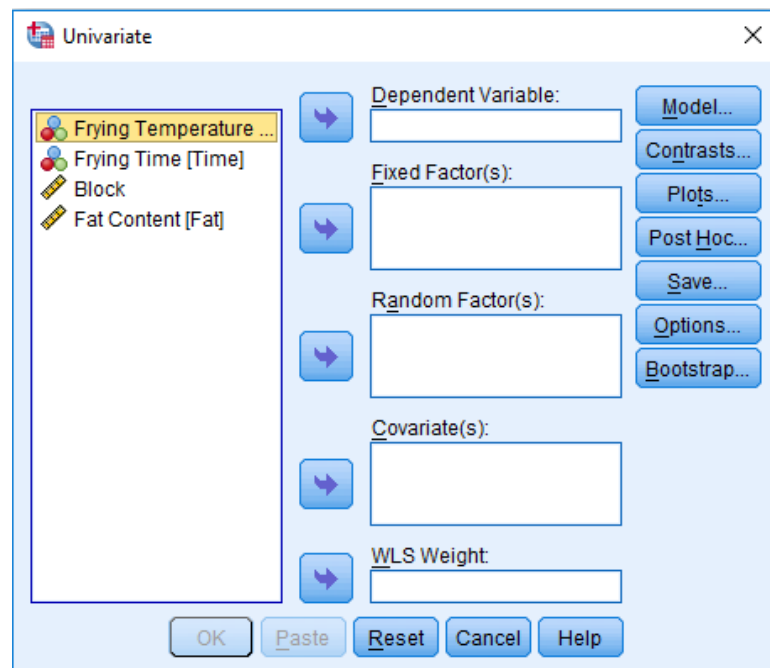
IBM SPSS Statistics Processor is ready

ภาพที่ 10.41 ป้อนข้อมูลตัวแปร Temp, Time, Block และ Fat ลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 10.42 จะปรากฏหน้าจอ ดังภาพที่ 10.43

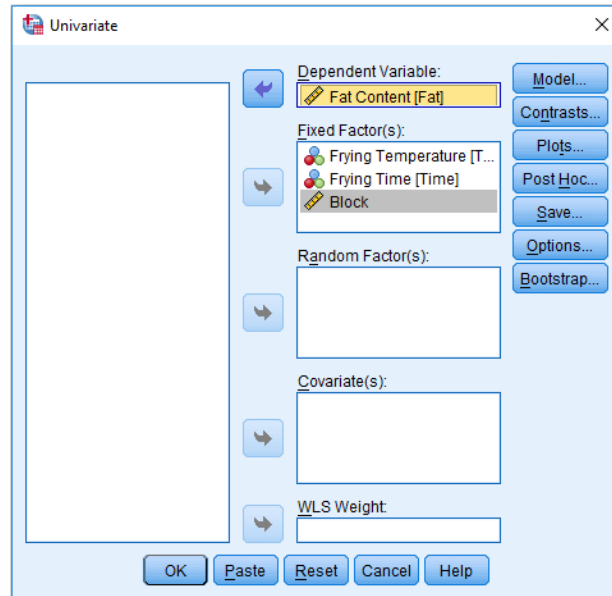


ภาพที่ 10.42 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



ภาพที่ 10.43 หน้าจอคำสั่ง Univariate

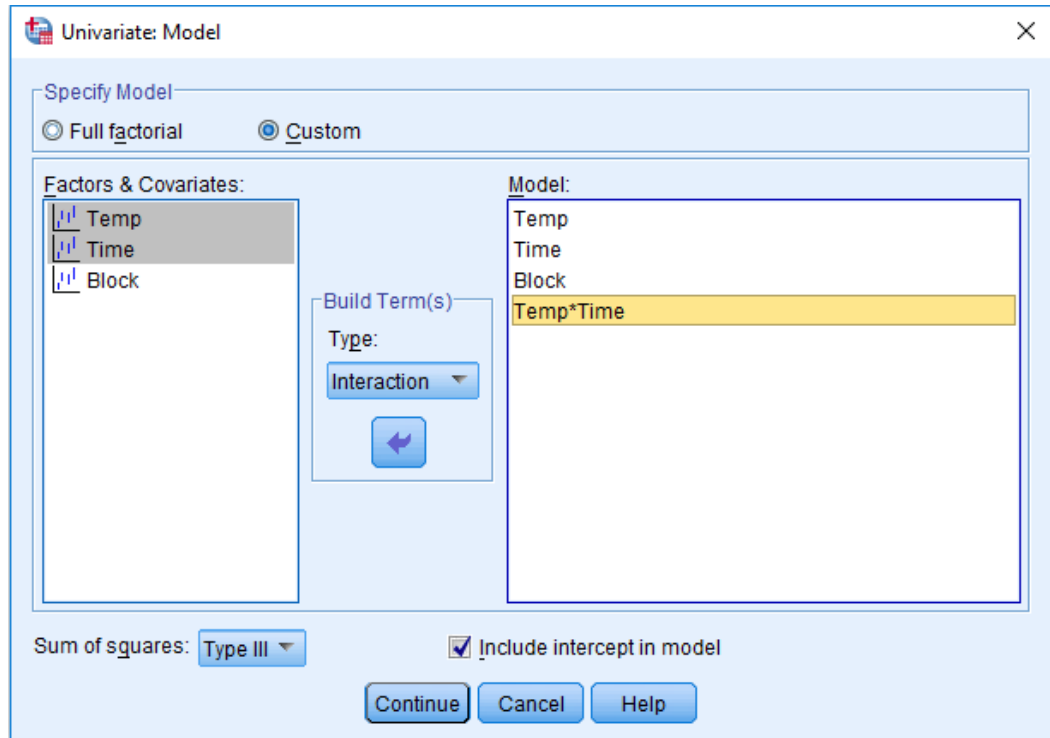
(4) เลือก ตัวแปร **Fat Content[Fat]** (ตัวแปรตาม) ใส่ในช่อง **Dependent Variable** และ ตัวแปร **Frying Temperature[Temp]** , ตัวแปร **Frying Time[Time]** และ ตัวแปร **Block** (ตัวแปรต้น) ใส่ในช่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.44



ภาพที่ 10.44 เลือกตัวแปร **Fat Content[Fat]** ใส่ในช่อง **Dependent Variable** และ เลือกตัวแปร **Frying Temperature[Temp]**, ตัวแปร **Frying Time[Time]** และตัวแปร **Block** ใส่ในช่อง **Fixed Factor(s)**

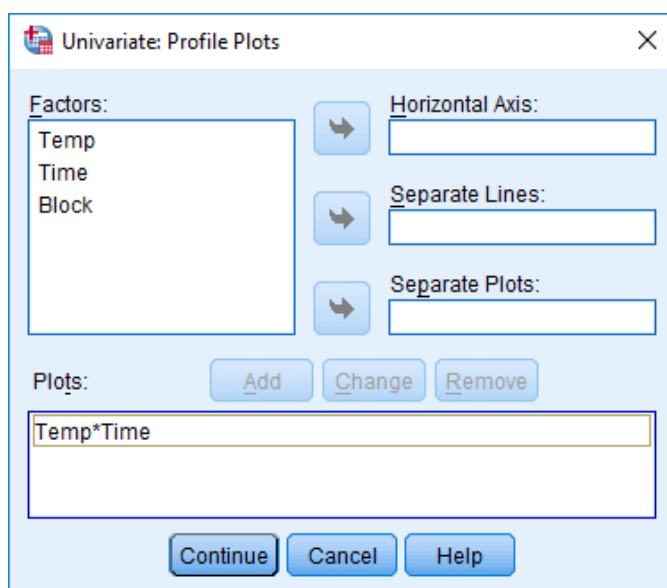
(5) เลือกคำสั่ง **Model...** จะได้หน้าจอตั้งภาพที่ 10.45 กำหนดค่าต่างๆ ดังนี้

- กำหนดรูปแบบการวิเคราะห์ (Model) ที่ตำแหน่ง Specify Model ให้คลิกเลือก **Custom** ทั้งนี้รูปแบบการวิเคราะห์ (Model) มี 2 รูปแบบ Full factorial คือ รูปแบบเต็มองค์ประกอบ และ Custom คือ การกำหนดรูปแบบขึ้นเอง
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลของปัจจัยหลัก (Main Effects) โดยกด Build Term(s) เลือก **Main Effects** แล้วเลือกตัวแปร **Temp, Time และ Block** ในกล่องซ้ายมือใส่ในช่อง Model ที่อยู่ในกล่องขวามือ
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลร่วมระหว่างปัจจัยทั้ง 2 ปัจจัย (Interaction Effects) โดยกด Build Term(s) เลือก **Interaction** แล้วกดคีย์บอร์ด **Shift ↑** ค้างไว้ กดเลือกตัวแปร Temp และ Time (กดเลือกแบบจับคู่) ในกล่องซ้ายมือใส่ในช่อง Model จะได้ตัวแปร **Temp*Time** ในกล่องขวามือ
- เมื่อกำหนดค่าเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



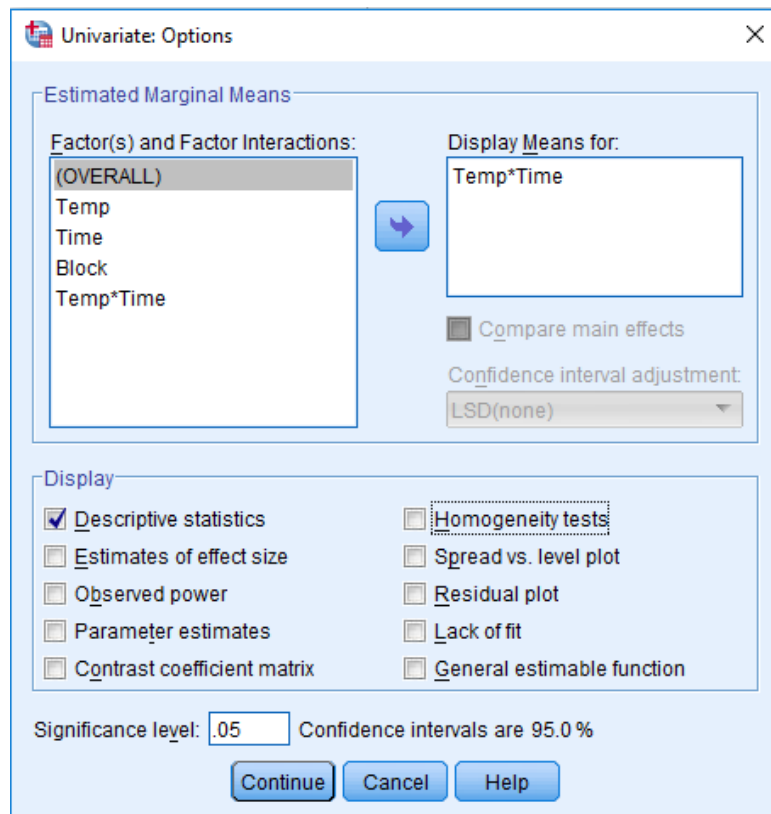
ภาพที่ 10.45 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลร่วมของทั้งสองปัจจัย จะได้หน้าจอตั้งภาพที่ 10.46 กดเลือกตัวแปร Temp ในช่อง Horizontal Axis (แกนนอน) และตัวแปร Time ในช่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Temp*Time** ในช่องด้านล่างเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 10.46 หน้าจอคำสั่ง Univariate : Profile Plots

(7) เลือกคำสั่ง **Options...** เพื่อให้แสดงค่าสถิติเบื้องต้นของแต่ละกลุ่ม จะได้หน้าจอ ดังภาพที่ 10.47 กดเลือก **ตัวแปร Temp*Time** ใส่ในช่อง Display Means for แล้วกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) เมื่อเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate และกด OK จะได้ผลลัพธ์จากการวิเคราะห์แสดงในหน้าจอ Output ดังภาพที่ 10.48



ภาพที่ 10.47 หน้าจอคำสั่ง Univariate : Options

Univariate Analysis of Variance

Between-Subjects Factors		
		N
Frying Temperature	150	6
	160	6
Frying Time	3	6
	5	6
Block	1	4
	2	4
	3	4

Descriptive Statistics

Dependent Variable: Fat Content

Frying Temperature	Frying Time	Block	Mean	Std. Deviation	N
150	3	1	27.600	.	1
		2	27.500	.	1
		3	27.600	.	1
		Total	27.567	.0577	3
	5	1	26.500	.	1
		2	26.600	.	1
		3	26.000	.	1
		Total	26.367	.3215	3
	Total	1	27.050	.7778	2
		2	27.050	.6364	2
		3	26.800	1.1314	2
		Total	26.967	.6890	6
160	3	1	25.700	.	1
		2	25.900	.	1
		3	26.100	.	1
		Total	25.900	.2000	3
	5	1	24.400	.	1
		2	24.600	.	1
		3	24.800	.	1
		Total	24.600	.2000	3
	Total	1	25.050	.9192	2
		2	25.250	.9192	2
		3	25.450	.9192	2
		Total	25.250	.7342	6
Total	3	1	26.650	1.3435	2
		2	26.700	1.1314	2
		3	26.850	1.0607	2
		Total	26.733	.9223	6
	5	1	25.450	1.4849	2
		2	25.600	1.4142	2
		3	25.400	.8485	2
		Total	25.483	.9968	6
	Total	1	26.050	1.3478	4
		2	26.150	1.2234	4
		3	26.125	1.1471	4
		Total	26.108	1.1245	12

Tests of Between-Subjects Effects

Dependent Variable: Fat Content

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	13.558 ^a	5	2.712	46.263	.000
Intercept	8179.741	1	8179.741	139559.559	.000
Temp	8.841	1	8.841	150.839	.000
Time	4.688	1	4.688	79.976	.000
Block	.022	2	.011	.185	.836
Temp * Time	.008	1	.008	.128	.733
Error	.352	6	.059		
Total	8193.650	12			
Corrected Total	13.909	11			

a. R Squared = .975 (Adjusted R Squared = .954)

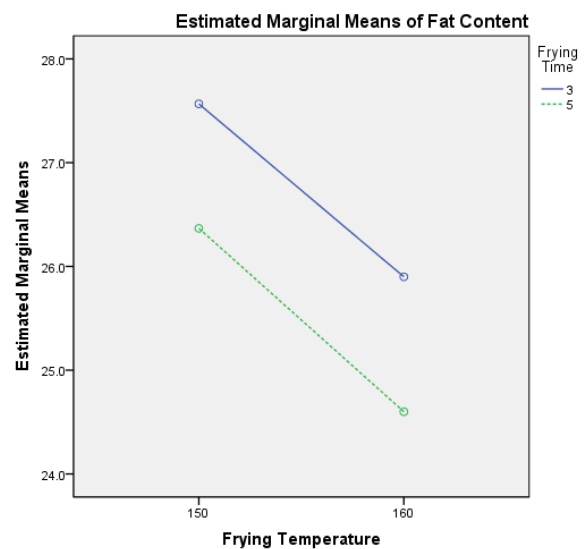
Estimated Marginal Means

Frying Temperature * Frying Time

Dependent Variable: Fat Content

Frying Temperature	Frying Time	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
150	3	27.567	.140	27.225	27.909
	5	26.367	.140	26.025	26.709
160	3	25.900	.140	25.558	26.242
	5	24.600	.140	24.258	24.942

Profile Plots



ภาพที่ 10.48 ผลลัพธ์จากการวิเคราะห์ความแปรปรวนข้อมูลการทดลองแบบ 2^2 Factorial in RCBD

10.6.4.4 การอ่านผลลัพธ์จากการวิเคราะห์ข้อมูลการทดลองแบบ 2^2 Factorial in RCBD

จากผลการวิเคราะห์ในภาพที่ 10.48 แสดงผลลัพธ์ได้ 5 ส่วน ดังนี้

ส่วนที่ 1 ตาราง Between-Subjects Factors เป็นส่วนที่แสดงรายละเอียดของข้อมูลที่ป้อน ดังนี้

Frying Temperature	คือ	คือ ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 150 และ 160 °C
Frying Time	คือ	คือ ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 3 และ 5 นาที
Block	คือ	คือ บล็อกหรือรอบในการทอด ในที่นี้มี 3 บล็อก ได้แก่ รอบที่ 1, 2 และ 3
N	คือ	คือ จำนวนข้อมูลในแต่ละกลุ่มที่จำแนกโดยตัวแปร Frying Temperature, ตัวแปร Frying Time และ ตัวแปร Block

ส่วนที่ 2 ตาราง Descriptive Statistics เป็นส่วนที่แสดงค่าสถิติเบื้องต้นของตัวแปรที่นำมาทดสอบ จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Frying Temperature	คือ	คือ ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 150 และ 160 °C
Frying Time	คือ	คือ ชื่อตัวแปรต้นที่ใช้จำแนกกลุ่ม ในที่นี้มี 2 กลุ่ม ได้แก่ 3 และ 5 นาที
Block	คือ	คือ บล็อกหรือรอบในการทอด ในที่นี้มี 3 บล็อก ได้แก่ รอบที่ 1, 2 และ 3
Mean	คือ	คือ ค่าเฉลี่ยปริมาณไขมันของแต่ละกลุ่ม
Std. Deviation	คือ	คือ ค่าส่วนเบี่ยงเบนมาตรฐานของปริมาณไขมัน ที่แสดงการกระจายของข้อมูล
N	คือ	คือ จำนวนข้อมูลในแต่ละกลุ่มที่จำแนกโดยตัวแปร Frying Temperature, ตัวแปร Frying Time และ ตัวแปร Block

ส่วนที่ 3 ตาราง Tests of Between-Subjects Effects เป็นส่วนที่แสดงค่าสถิติต่างๆ ของตาราง วิเคราะห์ความแปรปรวน เพื่อใช้ในการทดสอบสมมติฐานทางสถิติด้วยค่าสถิติ F หรือค่าความน่าจะเป็น Sig. เพื่อใช้ทดสอบสมมติฐาน อธิบายได้ ดังนี้

Type III Sum of Squares, df, Mean Square, F	คือ	คือ ค่าสถิติที่คำนวณได้จากข้อมูลตัวอย่าง
Sig.	คือ	คือ ค่าความน่าจะเป็นในการยอมรับหรือปฏิเสธสมมติฐาน H_0 ในที่นี้จะพิจารณาว่า Sig. ของตัวแปร Temp, ตัวแปร Time, ตัวแปร Temp*Time และ ตัวแปร Block ตามสมมติฐานที่เขียนไว้ ดังนี้

- (1) พิจารณาอิทธิพลของอุณหภูมิในการทอด พบว่าค่า Sig. ของตัวแปร Temp มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05
- (2) พิจารณาอิทธิพลของเวลาในการทอด พบว่าค่า Sig. ของตัวแปร Time มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05
- (3) พิจารณาอิทธิพลของปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอด พบว่าค่า Sig. ของตัวแปร Temp*Time มีค่าเท่ากับ 0.733 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันที่ระดับนัยสำคัญ 0.05
- (4) พิจารณาอิทธิพลของบล็อก พบว่าค่า Sig. ของตัวแปร Block มีค่าเท่ากับ 0.836 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

ส่วนที่ 4 ตาราง Estimated Marginal Means เป็นส่วนที่แสดงค่าเฉลี่ยโดยประมาณของตัวแปรที่นำมาทดสอบ เนื่องจากตอนวิเคราะห์เลือกตัวแปร Temp*Time ดังนั้นผลที่แสดงในตารางนี้จะเป็นค่าเฉลี่ยโดยประมาณของตัวแปร Frying Temperature แต่ละกลุ่มที่แยกตามระดับตัวแปร Frying Time จากข้อมูลดังกล่าว อธิบายได้ ดังนี้

Mean	คือ ค่าเฉลี่ยปริมาณไขมันโดยประมาณของแต่ละกลุ่มที่จำแนกด้วยตัวแปร Frying Temperature และตัวแปร Frying Time
Std. Error	คือ ค่าความคลาดเคลื่อนมาตรฐานในแต่ละกลุ่ม
95% Confidence Interval	คือ ค่าที่แสดงขอบเขตบนล่างของค่าเฉลี่ยในแต่ละกลุ่มในช่วงความเชื่อมั่น 95%

ส่วนที่ 5 Profile Plots เป็นส่วนที่แสดงกราฟเส้นของตัวแปรที่ต้องการทดสอบ

ในที่นี้ตัวแปรที่ต้องการทดสอบ คือ ตัวแปร Fat กราฟที่แสดงใช้ค่าเฉลี่ยโดยประมาณ (จากตาราง Estimated Marginal Means) มาแสดงโดยจำแนกตามตัวแปรที่ได้เลือกในขั้นตอนการวิเคราะห์ คำสั่ง Plots... ในที่นี้ได้เลือกตัวแปร Temp เป็นตัวแปรแกนนอน จึงเห็นกราฟในแกนนอนมี 2 สเกล (แสดงอุณหภูมิในการทอดที่ 150 และ 160 °C) และได้เลือกตัวแปร Time แสดงเป็นเส้นกราฟ จึงเห็นในกราฟมี 2 เส้น (แสดงเวลาในการทอดที่ 3 และ 5 นาที) จากกราฟพบว่ามีลักษณะเป็นเส้นขนานแสดงว่า 2 ปัจจัยนั้นไม่มี

อิทธิพลร่วมกัน ทั้งนี้เนื่องจากการนำค่าโดยประมาณมาใช้จึงอาจเกิดความผิดพลาดได้ ในการทดสอบอิทธิพลร่วมของทั้ง 2 ปัจจัยจึงควรพิจารณาจากค่า Sig. ในตาราง Tests of Between-Subjects Effects ก่อนเป็นอันดับแรกและใช้กราฟ Profile Plots ในการพิจารณาแนวโน้มค่าเฉลี่ยระหว่างระดับต่างๆ ของปัจจัยทั้ง 2 ตัว

10.6.4.5 การสรุปผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่กำหนดไว้ ดังนี้

สมมติฐานข้อที่ 1 การทดสอบสมมติฐานอิทธิพลของอุณหภูมิในการทอด (Main effect A)

H_0 : อุณหภูมิในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.48 ค่า Sig. ของตัวแปร Temp มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ อุณหภูมิในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 2 การทดสอบสมมติฐานอิทธิพลของเวลาในการทอด (Main effect B)

H_0 : เวลาในการทอดไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.48 ค่า Sig. ของตัวแปร Time มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ เวลาในการทอดมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 3 การทดสอบสมมติฐานอิทธิพลร่วมระหว่างอุณหภูมิในการทอดและเวลาในการทอด (Interaction effects AxB)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.49 ค่า Sig. ของตัวแปร Temp*Time มีค่าเท่ากับ 0.733 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ ไม่มีปฏิสัมพันธ์ระหว่างอุณหภูมิและเวลาในการทอดต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 4 การทดสอบสมมติฐานอิทธิพลของบล็อก (Block)

H_0 : บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

H_1 : บล็อกมีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

จากผลลัพธ์ที่ได้ในภาพที่ 10.48 ค่า Sig. ของตัวแปร Time มีค่าเท่ากับ 0.836 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ระดับนัยสำคัญ 0.05

10.6.4.6 การนำเสนอผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

การเขียนรายงานผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD ทำได้เช่นเดียวกับการวางแผนการทดลองแบบ 2^2 Factorial in CRD ในหัวข้อที่ 10.4.4.6 ซึ่งสามารถนำเสนอได้หลายแบบ ในที่นี้จะนำเสนอในรูปแบบตาราง 2 รูปแบบดังตารางที่ 10.13 และ 10.14 ซึ่งนักวิจัยจะเลือกนำเสนอแบบใดนั้นขึ้นอยู่กับวัตถุประสงค์ที่ต้องการนำเสนอข้อมูล

ตารางที่ 10.13 ค่า P-value แสดงอิทธิพลของอุณหภูมิในการทอด อิทธิพลของเวลาในการทอด อิทธิพลร่วมระหว่างอุณหภูมิและเวลาในการทอด และอิทธิพลของบล็อกต่อค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอด

อิทธิพลของปัจจัย	ค่า P-value
อุณหภูมิในการทอด	0.000*
เวลาในการทอด	0.000*
อุณหภูมิในการทอด × เวลาในการทอด	0.733
บล็อก	0.836

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.14 ค่าเฉลี่ยปริมาณไขมันในมันฝรั่งทอดที่ใช้อุณหภูมิและเวลาในการทอดต่างกัน

อุณหภูมิในการทอด (°C)	เวลาในการทอด (นาที)	ปริมาณไขมัน (%)
150	3	27.57 ± 0.06a
150	5	26.37 ± 0.32b
160	3	25.90 ± 0.20b
160	5	24.60 ± 0.20c
อิทธิพลของปัจจัย		ค่า P-value
อุณหภูมิในการทอด		0.000*
เวลาในการทอด		0.000*
อุณหภูมิในการทอด × เวลาในการทอด		0.733
บล็อก		0.836

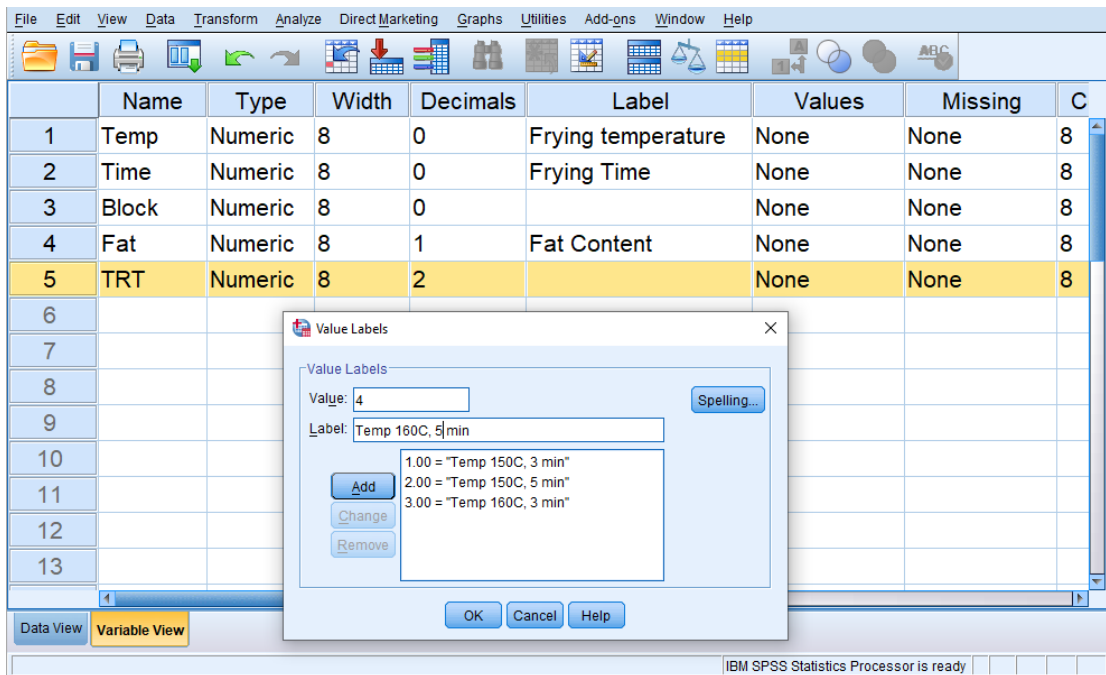
^{a-c} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ข้อสังเกต ในหนังสือเล่มนี้ ข้อมูลตัวอย่างที่ใช้ในการวิเคราะห์การวางแผนการทดลองแบบ 2^2 Factorial in CRD (ข้อ 10.4.4) และ 2^2 Factorial in RCBD (ข้อ 10.6.4) เป็นข้อมูลชุดเดียวกัน เมื่อสังเกตผลการวิเคราะห์ (ตารางที่ 10.7 และ 10.13) จะพบว่าค่าเฉลี่ยปริมาณไขมันในแต่ละสิ่งทดลองเท่ากันแม้ว่าจะใช้แผนการทดลองต่างกัน แต่สิ่งที่แตกต่างกัน คือ (1) ค่า Sig. หรือค่า P-Value ที่ได้จากการวิเคราะห์ความแปรปรวน และ (2) ความแตกต่างของค่าเฉลี่ยแต่ละสิ่งทดลองจากการวิเคราะห์ Post Hoc Tests เนื่องจากในการวางแผนการทดลองแบบ Factorial in RCBD ความแปรปรวนบางส่วนเกิดจากบล็อก ดังนั้นนักวิจัยควรเลือกใช้การวางแผนการทดลองให้เหมาะสมกับลักษณะงานวิจัยของตนเอง โดยการพิจารณาความสม่ำเสมอของหน่วยทดลอง หากหน่วยทดลองมีความสม่ำเสมอให้เลือกใช้แผนการทดลองแบบ Factorial in CRD กรณีที่ไม่แน่ใจว่าหน่วยทดลองมีความสม่ำเสมอหรือไม่ ให้ใช้การวางแผนการทดลองแบบ Factorial in RCBD

ตารางที่ 10.14 ได้มาจากการกำหนดตัวแปร var ใหม่ชื่อ TRT ซึ่งเป็นตัวแปรที่แสดงถึงสิ่งทดลองที่เกิดจากปัจจัยร่วมระหว่างอุณหภูมิในการทอด (2 ระดับ) และ เวลาในการทอด(2 ระดับ) หรือเรียกว่า **“Treatment Combination”** ในตัวอย่างนี้จะได้จำนวนสิ่งทดลอง (TRT) เท่ากับ 4 สิ่งทดลอง (คำนวณจากผลคูณของระดับแต่ละปัจจัย) นำข้อมูลตัวแปร TRT, ตัวแปร Block และตัวแปร Fat ไปวิเคราะห์ด้วยคำสั่ง **Post Hoc...** เพื่อเปรียบเทียบค่าเฉลี่ยแบบจับคู่ทุกคู่ ทำเช่นเดียวกับการวิเคราะห์ความแปรปรวนของข้อมูลที่ได้จากการวางแผนการทดลองแบบ RCBD (หัวข้อที่ 9.6.3) ตามขั้นตอนดังนี้

(1) กำหนดรายละเอียดและค่า Values ของตัวแปร TRT ในหน้าต่าง Variable View ดังภาพที่ 10.49 และป้อนข้อมูล TRT ในหน้าต่าง Data View ดังภาพที่ 10.50

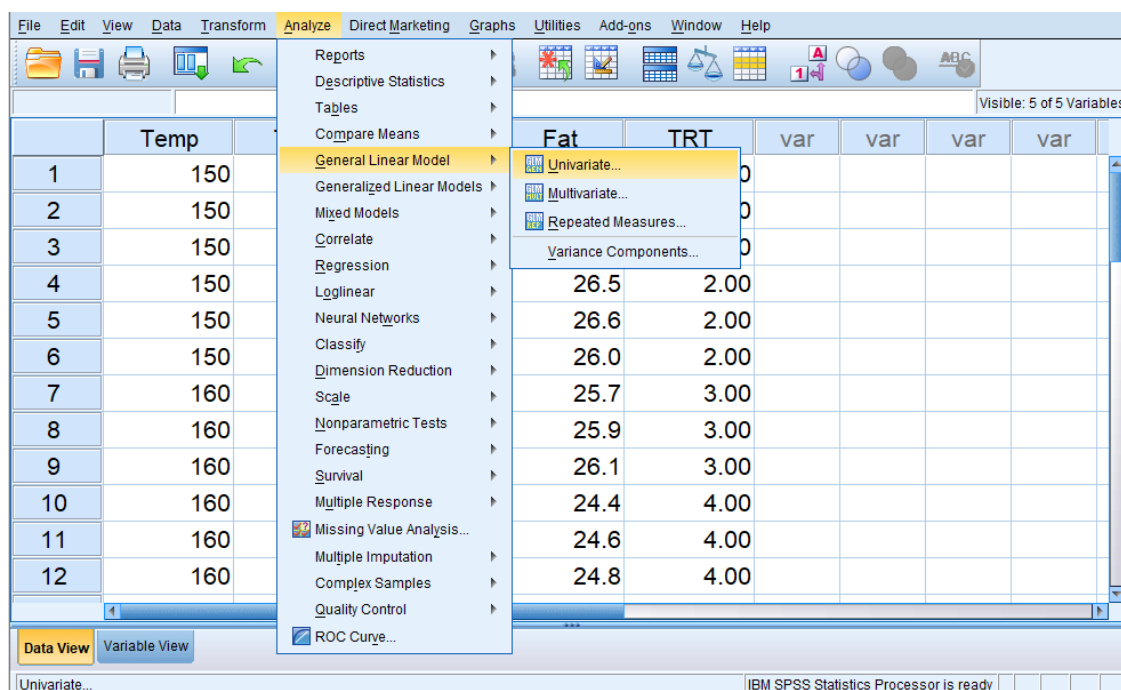


ภาพที่ 10.49 การกำหนดรายละเอียดตัวแปร TRT ในหน้าต่าง Variable View

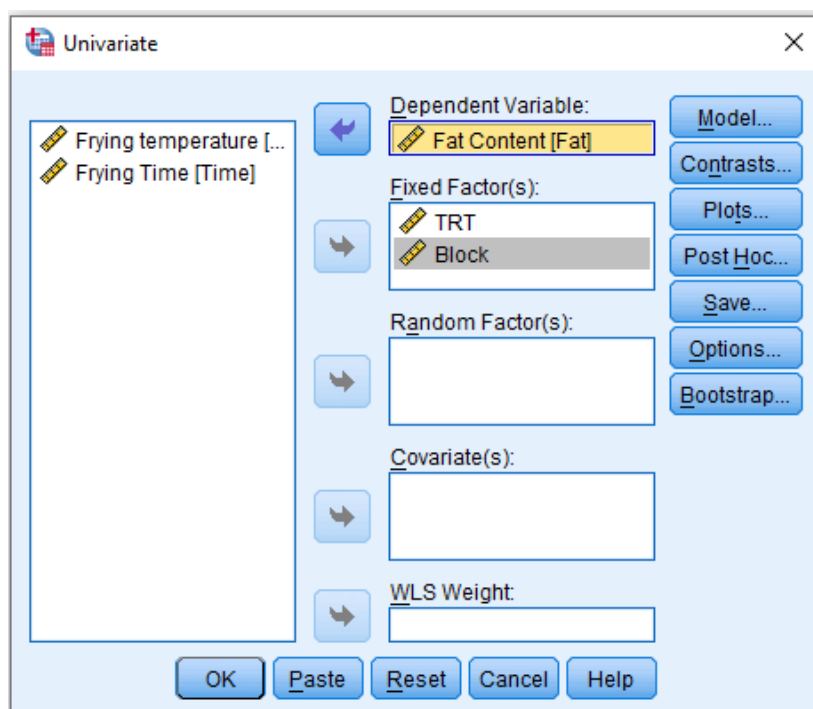
	Temp	Time	Block	Fat	TRT	var	var	var	var
1	150	3	1	27.6	1.00				
2	150	3	2	27.5	1.00				
3	150	3	3	27.6	1.00				
4	150	5	1	26.5	2.00				
5	150	5	2	26.6	2.00				
6	150	5	3	26.0	2.00				
7	160	3	1	25.7	3.00				
8	160	3	2	25.9	3.00				
9	160	3	3	26.1	3.00				
10	160	5	1	24.4	4.00				
11	160	5	2	24.6	4.00				
12	160	5	3	24.8	4.00				

ภาพที่ 10.50 การป้อนข้อมูล TRT ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 10.51 และเลือกตัวแปร Fat Content [Fat] (ตัวแปรตาม) ใส่ในกล่อง **Dependent Variable** และ ตัวแปร TRT กับ ตัวแปร Block ใส่ในกล่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.52

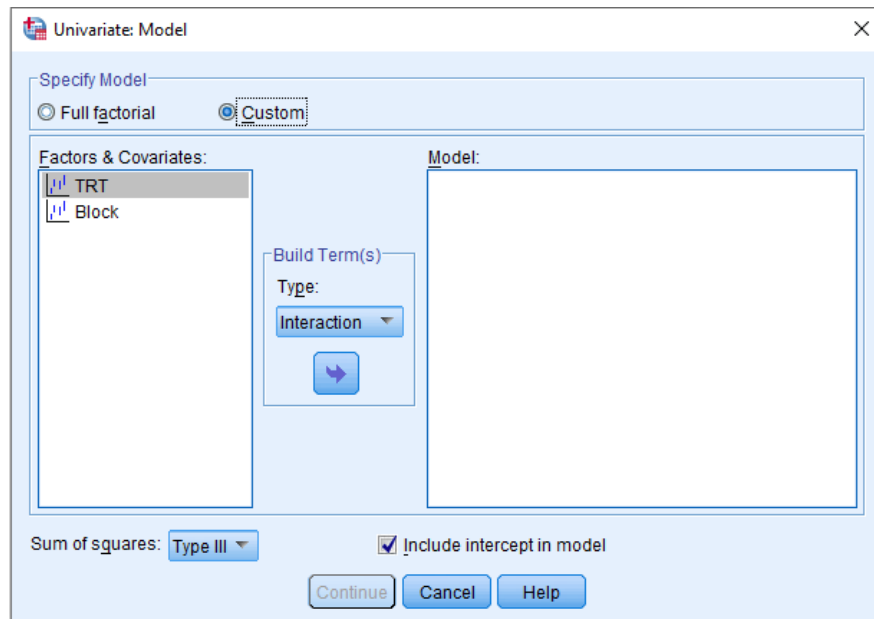


ภาพที่ 10.51 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...

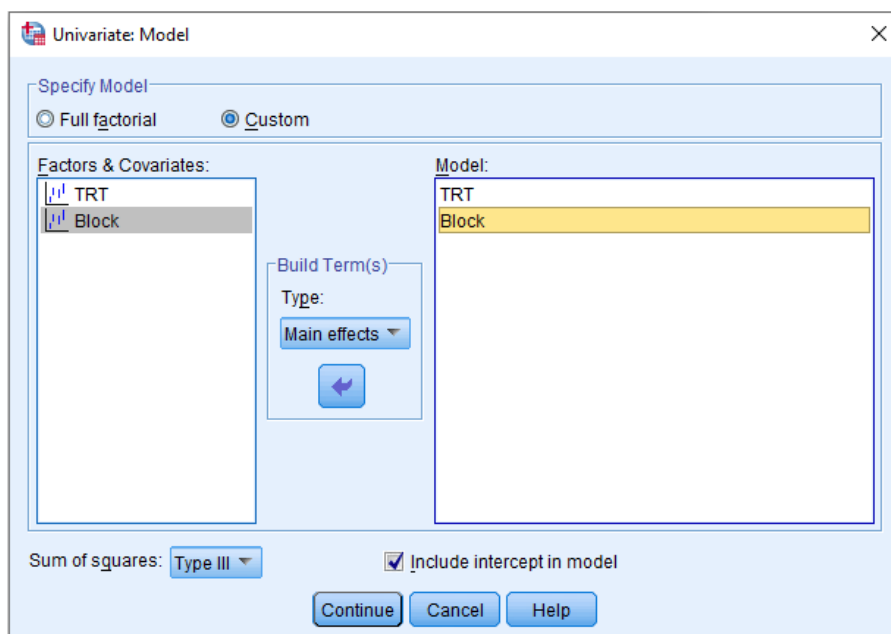


ภาพที่ 10.52 การกำหนดตัวแปรในเมนูคำสั่ง Univariate

(3) เลือกคำสั่ง **Model...** ในการวางแผนการทดลองแบบ Factorial in RCBD ต้องกำหนดตัวแบบ (Model) ชนิด **Custom** (ดังภาพที่ 10.53) จากนั้นให้กดเลือก **Type** ของ **Build Terms(s)** เป็น **Main effects** แล้วกดเลือก **ตัวแปร TRT** และ **ตัวแปร Block** จากกล่องซ้ายมือ (Factors & Covariates) ไปใส่ในกล่องขวามือ (Model) โดยคลิกเลือกทีละตัวแปร (ดังภาพที่ 10.54) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate (ข้อสังเกต จะไม่กำหนด Build Term(s) แบบ Interaction เนื่องจากเงื่อนไขของแผนการทดลองแบบ RCBD ต้องไม่มีปฏิสัมพันธ์ระหว่างสิ่งทดลองและบล็อก)

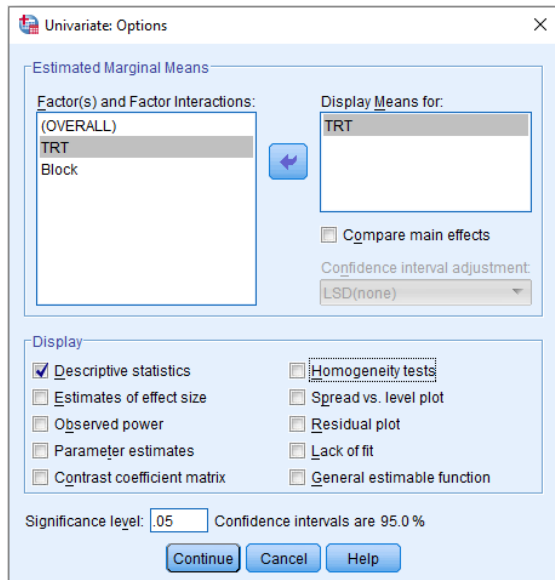


ภาพที่ 10.53 หน้าจอคำสั่ง Univariate : Model



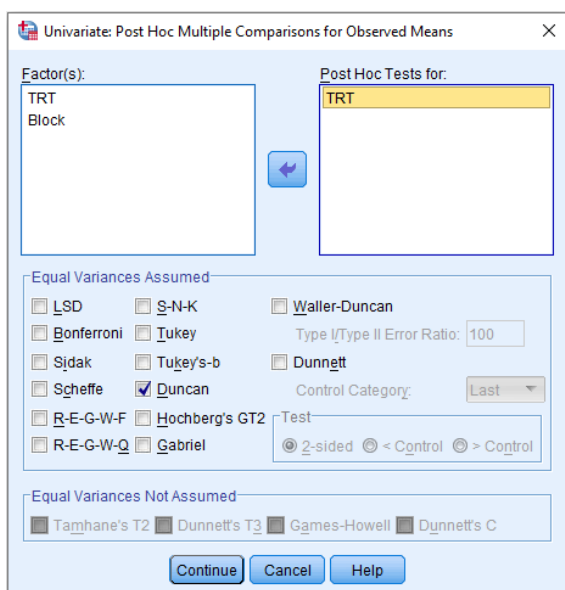
ภาพที่ 10.54 การเลือกตัวแปรในกล่อง Model

(4) เลือกคำสั่ง **Options...** จะได้หน้าจอตั้งภาพที่ 10.55 ให้เลือก **ตัวแปร TRT** จากกล่องซ้ายมือ (Factor(s) and Factor Interactions) ไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 10.55 หน้าจอคำสั่ง Univariate : Options

(5) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือ เปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอตั้งภาพที่ 10.56 ให้กดเลือก **ตัวแปร TRT** จากกล่องซ้ายมือ (Factor(s)) ไปใส่ในกล่องขวามือ (Post Hoc Tests for) และกดเลือก **วิธี Duncan** (นักวิจัยสามารถเลือกวิธีอื่นที่ต้องการใช้ทดสอบได้หลายวิธี) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 10.57



ภาพที่ 10.56 หน้าจอคำสั่ง Univariate : Post Hoc Multiple Comparisons for Observed Means

Univariate Analysis of Variance

Between-Subjects Factors			
	Value Label	N	
TRT	1.00 Temp 150C, 3 min	3	
	2.00 Temp 150C, 5 min	3	
	3.00 Temp 160C, 3 min	3	
	4.00 Temp 160C, 5 min	3	
Block	1	4	
	2	4	
	3	4	

Descriptive Statistics				
Dependent Variable: Fat Content				
TRT	Block	Mean	Std. Deviation	N
Temp 150C, 3 min	1	27.600	.	1
	2	27.500	.	1
	3	27.600	.	1
	Total	27.567	.0577	3
Temp 150C, 5 min	1	26.500	.	1
	2	26.600	.	1
	3	26.000	.	1
	Total	26.367	.3215	3
Temp 160C, 3 min	1	25.700	.	1
	2	25.900	.	1
	3	26.100	.	1
	Total	25.900	.2000	3
Temp 160C, 5 min	1	24.400	.	1
	2	24.600	.	1
	3	24.800	.	1
	Total	24.600	.2000	3
Total	1	26.050	1.3478	4
	2	26.150	1.2234	4
	3	26.125	1.1471	4
	Total	26.108	1.1245	12

Tests of Between-Subjects Effects					
Dependent Variable: Fat Content					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	13.558 ^a	5	2.712	46.263	.000
Intercept	8179.741	1	8179.741	139559.559	.000
TRT	13.536	3	4.512	76.981	.000
Block	.022	2	.011	.185	.836
Error	.352	6	.059		
Total	8193.650	12			
Corrected Total	13.909	11			

a. R Squared = .975 (Adjusted R Squared = .954)

Estimated Marginal Means

TRT				
Dependent Variable: Fat Content				
TRT	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
Temp 150C, 3 min	27.567	.140	27.225	27.909
Temp 150C, 5 min	26.367	.140	26.025	26.709
Temp 160C, 3 min	25.900	.140	25.558	26.242
Temp 160C, 5 min	24.600	.140	24.258	24.942

Post Hoc Tests

TRT

Homogeneous Subsets

Fat Content				
Duncan ^{a,b}				
TRT	N	Subset		
		1	2	3
Temp 160C, 5 min	3	24.600		
Temp 160C, 3 min	3		25.900	
Temp 150C, 5 min	3		26.367	
Temp 150C, 3 min	3			27.567
Sig.		1.000	.056	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .059.

a. Uses Harmonic Mean Sample Size = 3.000.
b. Alpha = .05.

ภาพที่ 10.57 ผลลัพธ์การวิเคราะห์ความแปรปรวนของตัวแปร TRT และ ตัวแปร Fat

10.7 กรณีศึกษาการวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

10.7.1 ตัวอย่างข้อมูล

นักวิจัยได้ทำการสำรวจตลาดแล้วพบว่าผลิตภัณฑ์ที่มีส่วนผสมของทุเรียนกำลังเป็นที่ต้องการของตลาด ดังนั้นนักวิจัยจึงต้องการพัฒนาสูตรเค้กทุเรียนโดยใช้เนื้อทุเรียนแทนเนื้อมะพร้าวในสูตรเค้กพื้นฐาน แต่เนื่องจากเนื้อทุเรียนสุกมีความหวานมากหากใช้ปริมาณน้ำตาลเท่าเดิมอาจส่งผลกระทบต่อกรยอมรับของผู้บริโภค นักวิจัยจึงศึกษาปริมาณเนื้อทุเรียนสุกบด 2 ระดับ (30 และ 40 กรัม) และปริมาณน้ำตาล 2 ระดับ (100 และ 150 กรัม) และควบคุมให้ส่วนผสมอื่นคงที่ให้ผู้ทดสอบชิมจำนวน 30 คนประเมินความชอบด้วยสเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด) วางแผนการทดลองแบบ 2^2 Factorial in RCBD โดยกำหนดให้บล็อก คือ ผู้ทดสอบชิม ได้ข้อมูลดังตารางที่ 10.15

ตารางที่ 10.15 คะแนนความชอบเค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียนและปริมาณน้ำตาลแตกต่างกัน

ผู้ชิม คนที่	สูตรเค้กทุเรียน			
	สูตร 1	สูตร 2	สูตร 3	สูตร 4
1	6	7	6	5
2	6	7	6	5
3	6	7	6	5
4	6	7	6	5
5	6	7	6	5
6	6	8	6	6
7	6	8	6	6
8	7	8	6	6
9	7	8	6	6
10	7	7	7	6
11	7	7	7	6
12	7	7	7	5
13	7	7	7	5
14	7	6	7	5
15	7	6	7	6

ผู้ชิม คนที่	สูตรเค้กทุเรียน			
	สูตร 1	สูตร 2	สูตร 3	สูตร 4
16	6	6	7	6
17	6	6	7	6
18	6	6	7	6
19	6	8	7	5
20	6	8	7	5
21	6	8	6	5
22	7	7	6	6
23	7	7	6	6
24	7	7	6	6
25	7	6	6	6
26	7	6	6	6
27	7	6	6	5
28	6	6	6	5
29	6	6	6	5
30	6	7	6	5

สูตร 1 : ใช้เนื้อทุเรียน 30 กรัม น้ำตาล 100 กรัม, สูตร 2 : ใช้เนื้อทุเรียน 30 กรัม น้ำตาล 150 กรัม

สูตร 3 : ใช้เนื้อทุเรียน 40 กรัม น้ำตาล 100 กรัม, สูตร 4 : ใช้เนื้อทุเรียน 40 กรัม น้ำตาล 150 กรัม

จากข้อมูลดังกล่าว พิจารณารายละเอียดของข้อมูลได้ ดังนี้

- ปัจจัย A คือ ปริมาณเนื้อทุเรียน มี 2 ระดับ (30 และ 40 กรัม)
- ปัจจัย B คือ ปริมาณน้ำตาล มี 2 ระดับ (100 และ 150 กรัม)
- บล็อก คือ ผู้ทดสอบชิม มี 30 บล็อก (จำนวนผู้ทดสอบชิม)
- จำนวนสิ่งทดลอง (Treatment) เท่ากับ 4 สิ่งทดลอง (ผลคูณของระดับแต่ละปัจจัย)
- จำนวนหน่วยทดลอง (Experimental unit) ในแต่ละสิ่งทดลอง เท่ากับ 30 หน่วย
- จำนวนหน่วยทดลอง (Experimental unit) ทั้งหมด เท่ากับ 120 หน่วย
- ค่าสังเกต หรือค่าตอบสนอง (ตัวแปรตาม) คือ ค่าคะแนนความชอบ

10.7.2 การเขียนสมมติฐานเพื่อวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

สามารถเขียนสมมติฐานสำหรับทดสอบได้ ดังนี้

สมมติฐานข้อที่ 1 ทดสอบอิทธิพลของระดับปริมาณเนื้อทุเรียน (Main effect A)

- | | | |
|-------|---|---|
| H_0 | : | เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30 และ 40 กรัม มีค่าเฉลี่ยคะแนนความชอบไม่แตกต่างกัน |
| H_1 | : | เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30 และ 40 กรัม มีค่าเฉลี่ยคะแนนความชอบแตกต่างกัน |

สมมติฐานข้อที่ 2 ทดสอบอิทธิพลของระดับปริมาณน้ำตาล (Main effect B)

- | | | |
|-------|---|---|
| H_0 | : | เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยคะแนนความชอบไม่แตกต่างกัน |
| H_1 | : | เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยคะแนนความชอบแตกต่างกัน |

สมมติฐานข้อที่ 3 ทดสอบอิทธิพลร่วมระหว่างระดับปริมาณเนื้อทุเรียนและระดับปริมาณน้ำตาล (Interaction effects AxB)

- | | | |
|-------|---|--|
| H_0 | : | ไม่มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยคะแนนความชอบ |
| H_1 | : | มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยคะแนนความชอบ |

สมมติฐานข้อที่ 4 การทดสอบสมมติฐานอิทธิพลของบล็อก (Block)

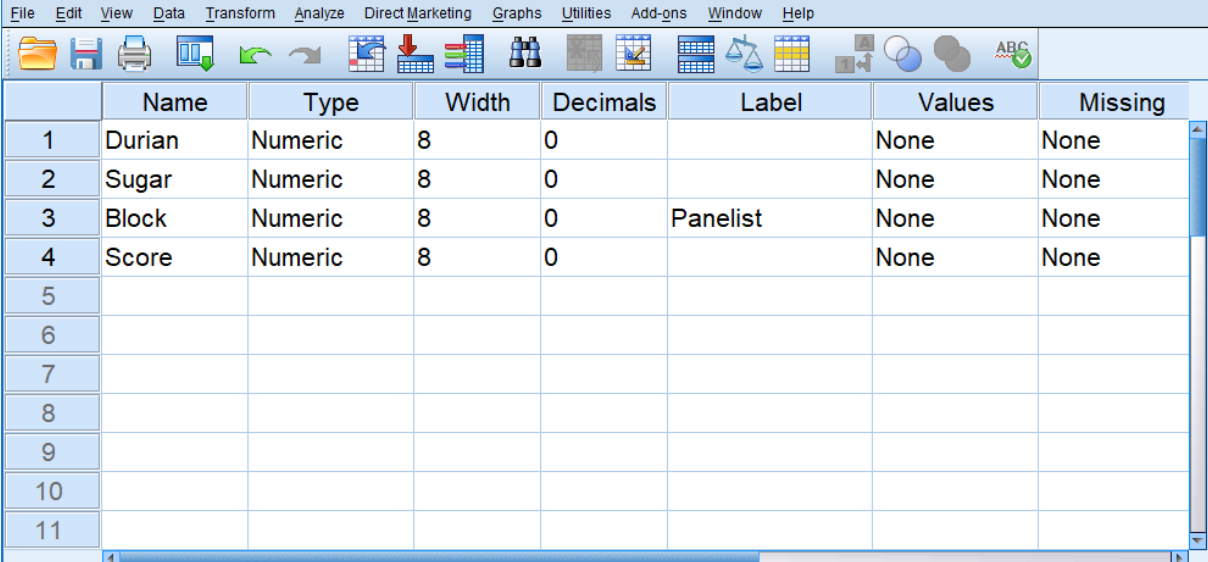
- | | | |
|-------|---|---|
| H_0 | : | บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบ |
| H_1 | : | บล็อกมีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบ |

10.7.3 การใช้คำสั่งใน SPSS วิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 2² Factorial in RCBD

(1) เปิดโปรแกรม SPSS ในการป้อนข้อมูลค่าคะแนนความชอบของเค้กทุเรียนแต่ละสิ่งทดลอง จะต้องใช้ตัวแปร var จำนวน 4 ตัวแปรสำหรับป้อนข้อมูล 4 ตัว ได้แก่ ปริมาณเนื้อทุเรียน, ปริมาณน้ำตาล, ลำดับผู้ทดสอบ และ ค่าคะแนนความชอบ ดังนั้นกำหนดชื่อตัวแปรเป็น **Durian, Sugar, Block** และ **Score** ในหน้าต่าง Variable View ดังภาพที่ 10.58 โดยกำหนดชื่อตัวแปร (Name), ความหมายที่แท้จริงของตัวแปร (Label), จำนวนทศนิยม (Decimals) และค่ารหัส (Values) ดังนี้

ข้อมูลที่ต้องการป้อน	กำหนด Name	กำหนด Label	กำหนด Decimals	กำหนด Values
ระดับปริมาณเนื้อทุเรียน	Durian	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด
ระดับปริมาณน้ำตาล	Sugar	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด
ลำดับผู้ทดสอบ	Block	Panelist	0	ไม่ต้องกำหนด
ค่าคะแนนความชอบ	Score	ไม่ต้องกำหนด	0	ไม่ต้องกำหนด

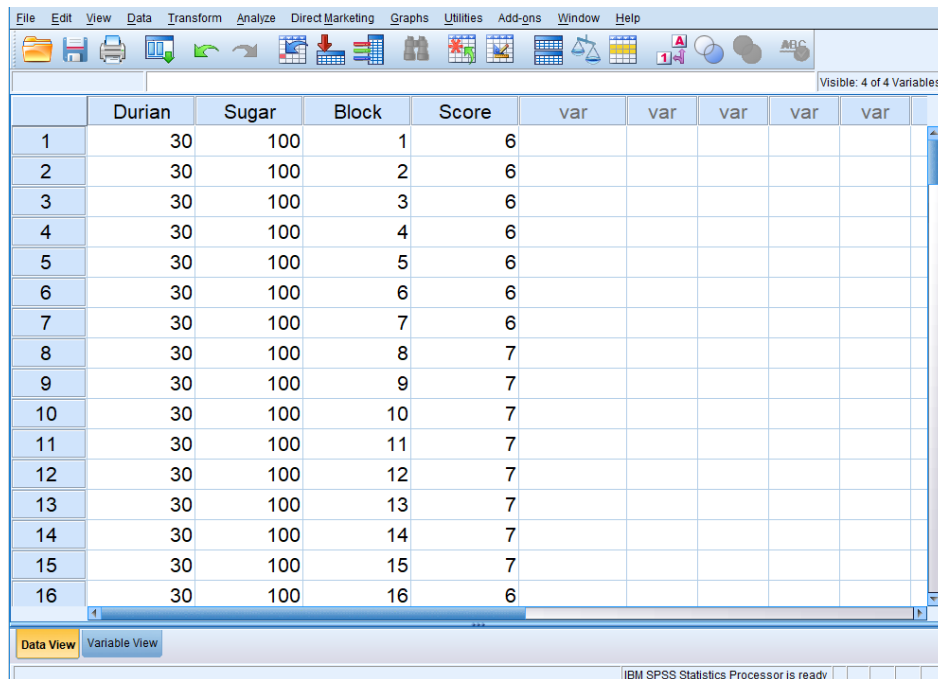
หมายเหตุ กรณีนี้กวีจายำชื่อตัวแปร (Name) ได้ว่ามาจากข้อมูลใด ไม่จำเป็นต้องกำหนดคำอธิบาย Label ได้



	Name	Type	Width	Decimals	Label	Values	Missing
1	Durian	Numeric	8	0		None	None
2	Sugar	Numeric	8	0		None	None
3	Block	Numeric	8	0	Panelist	None	None
4	Score	Numeric	8	0		None	None
5							
6							
7							
8							
9							
10							
11							

ภาพที่ 10.58 กำหนดรายละเอียดตัวแปร Durian, Sugar, Block และ Score ในหน้าต่าง Variable View

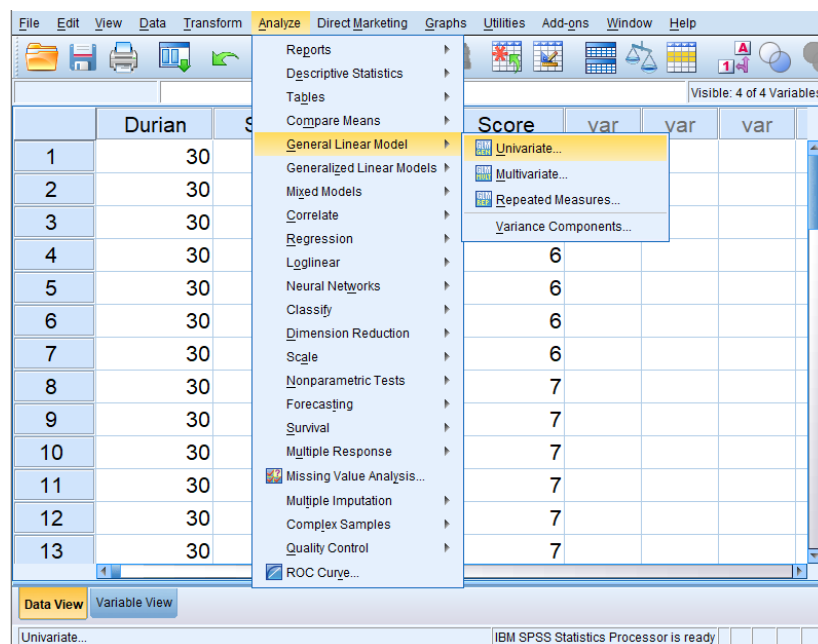
(2) ป้อนข้อมูลลงในหน้าต่าง Data View ดังภาพที่ 10.59 โดยป้อนข้อมูลค่าคะแนนความชอบของเค้กทุเรียนแต่ละกลุ่มเรียงตามระดับปริมาณเนื้อทุเรียนและปริมาณน้ำตาลและตามลำดับผู้ทดสอบชิม



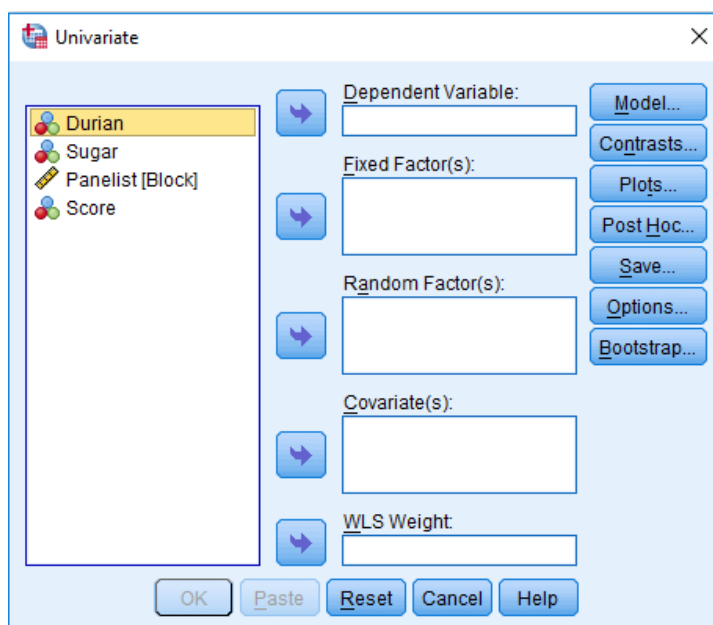
	Durian	Sugar	Block	Score	var	var	var	var	var
1	30	100	1	6					
2	30	100	2	6					
3	30	100	3	6					
4	30	100	4	6					
5	30	100	5	6					
6	30	100	6	6					
7	30	100	7	6					
8	30	100	8	7					
9	30	100	9	7					
10	30	100	10	7					
11	30	100	11	7					
12	30	100	12	7					
13	30	100	13	7					
14	30	100	14	7					
15	30	100	15	7					
16	30	100	16	6					

ภาพที่ 10.59 ป้อนข้อมูลลงในหน้าต่าง Data View

(3) เลือกเมนู **Analyze > General Linear Model > Univariate...** ดังภาพที่ 10.60 จะปรากฏหน้าต่างดังภาพที่ 10.61 จะเห็นว่าในกล่องซ้ายมือจะแสดงชื่อของตัวแปรแต่ละตัวที่กำหนดไว้

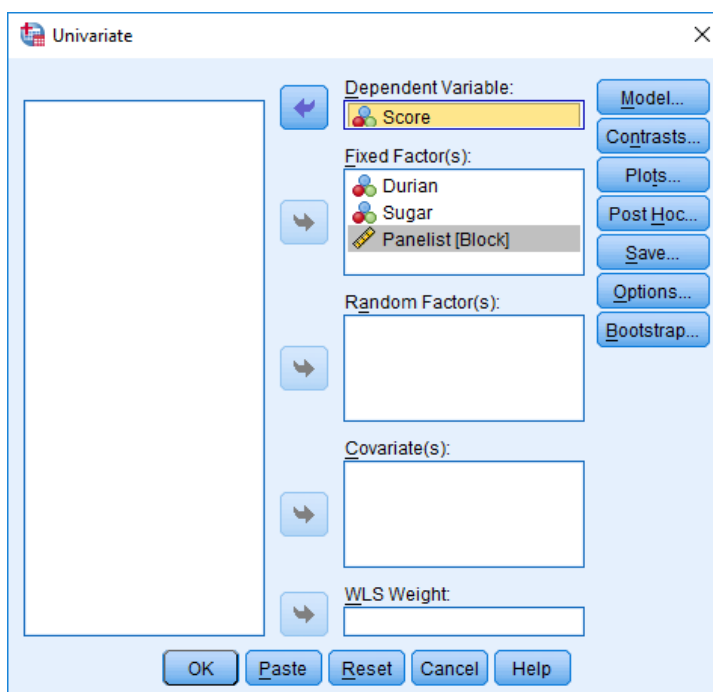


ภาพที่ 10.60 เมนูหน้าจอคำสั่ง Analyze > General Linear Model > Univariate...



ภาพที่ 10.61 หน้าจोकำสั่ง Univariate

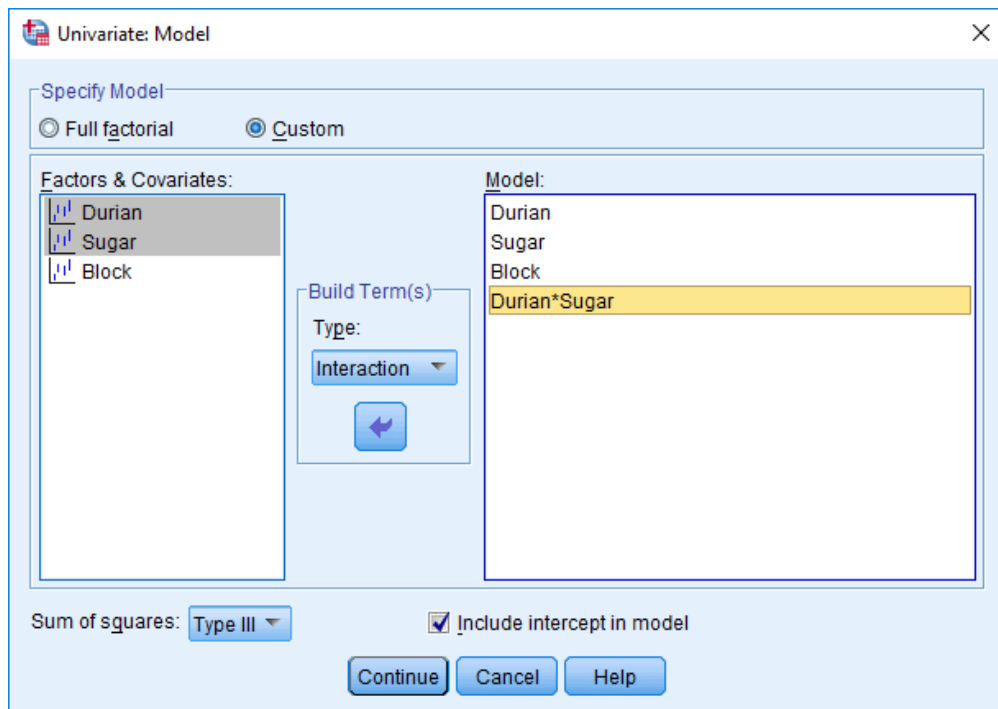
(4) เลือก ตัวแปร Score (ตัวแปรตาม) ใส่ในช่อง **Dependent Variable** และ ตัวแปร Durian, ตัวแปร Sugar และ ตัวแปร Panelist [Block] ใส่ในช่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.62



ภาพที่ 10.62 เลือกตัวแปร Score ใส่ในช่อง Dependent Variable และ เลือกตัวแปร Durian, ตัวแปร Sugar และตัวแปร Panelist [Block] ใส่ในช่อง Fixed Factor(s)

(5) เลือกคำสั่ง **Model...** จะได้หน้าจอตั้งภาพที่ 10.63 กำหนดค่าต่างๆ ดังนี้

- กำหนดรูปแบบการวิเคราะห์ (Model) ที่ตำแหน่ง Specify Model ให้คลิกเลือก **Custom** ทั้งนี้รูปแบบการวิเคราะห์ (Model) มี 2 รูปแบบ Full factorial คือ รูปแบบเต็มองค์ประกอบ และ Custom คือ การกำหนดรูปแบบขึ้นเอง
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลของปัจจัยหลัก (Main Effects) โดยกด Build Term(s) เลือก **Main Effects** แล้วเลือก ตัวแปร Durian, ตัวแปร Sugar และ ตัวแปร Block ในกล่องซ้ายมือใส่ในช่อง Model ที่อยู่ในกล่องขวามือ
- กำหนดตัวแปรเพื่อทดสอบอิทธิพลร่วมระหว่างปัจจัยทั้ง 2 ปัจจัย (Interaction Effects) โดยกด Build Term(s) เลือก **Interaction** แล้วกดคีย์บอร์ด **Shift ↑** ค้างไว้ กดเลือกตัวแปร Durian และ ตัวแปร Sugar (กดเลือกแบบจับคู่) ในกล่องซ้ายมือใส่ในช่อง Model จะได้ตัวแปร **Durian*Sugar** ในกล่องขวามือ
- เมื่อกำหนดค่าเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



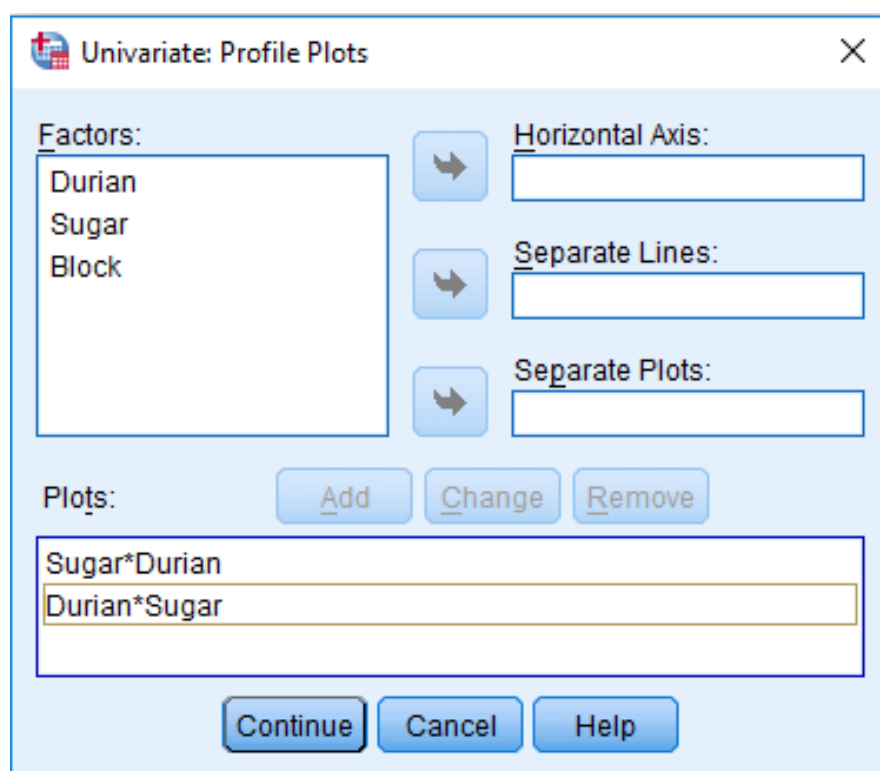
ภาพที่ 10.63 หน้าจอคำสั่ง Univariate : Model

(6) เลือกคำสั่ง **Plots...** เพื่อพิจารณาอิทธิพลร่วมของทั้งสองปัจจัย จะได้หน้าจอตั้งภาพที่ 10.64 ในตัวอย่างนี้จะสร้าง 2 กราฟ โดยสลับตัวแปรที่ใช้เป็นแกนนอนและที่ใช้แสดงรูปแบบเส้นกราฟ

- กราฟที่ 1 กดเลือกตัวแปร Sugar ในช่อง Horizontal Axis (แกนนอน) และตัวแปร Durian ในช่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Sugar*Durian** ในช่องด้านล่าง

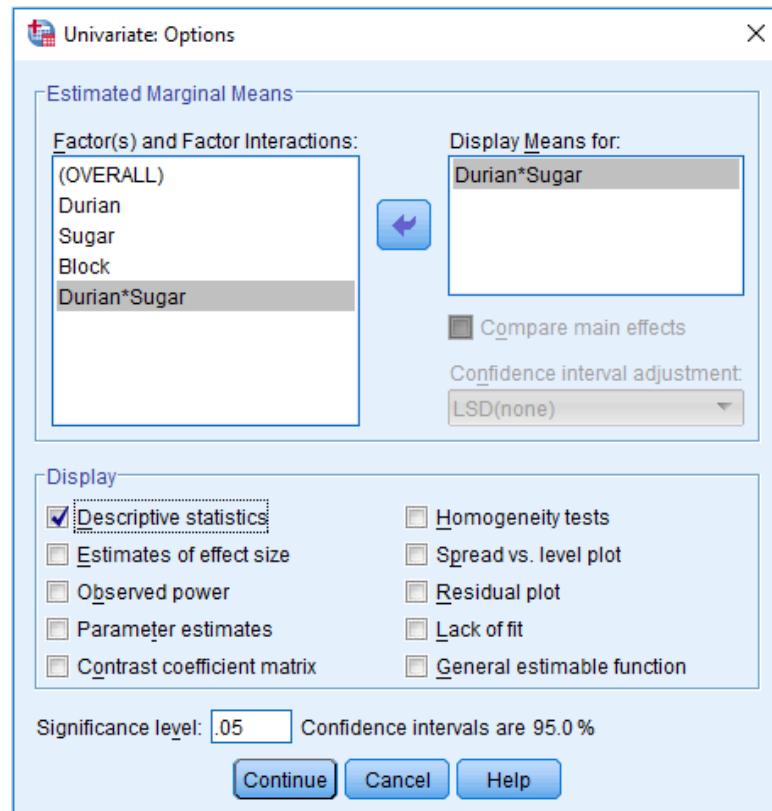
- กราฟที่ 2 กดเลือกตัวแปร Durian ในช่อง Horizontal Axis (แกนนอน) และตัวแปร Sugar ในช่อง Separate Lines (แสดงข้อมูลในรูปแบบเส้นกราฟ) ต่อจากนั้นให้กด Add จะได้ตัวแปร **Durian*Sugar** ในช่องด้านล่าง

- เมื่อกำหนดตัวแปรเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate



ภาพที่ 10.64 หน้าจอคำสั่ง Univariate : Profile Plots

(7) เลือกคำสั่ง **Options...** เพื่อให้แสดงค่าสถิติเบื้องต้นของแต่ละกลุ่ม จะได้หน้าจอตั้งภาพที่ 10.65 กดเลือก **ตัวแปร Durian*Sugar** ใส่ในช่อง Display Means for แล้วกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) เมื่อเสร็จแล้วให้กด Continue เพื่อกลับไปหน้าจอหลัก Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 10.66



ภาพที่ 10.65 หน้าจอคำสั่ง Univariate : Options

Univariate Analysis of Variance

Between-Subjects Factors

	N
Durian 30	60
40	60
Sugar 100	60
150	60
Panelist 1	4
2	4
3	4
4	4
5	4
6	4

Rows 1 through 10 of 34

Descriptive Statistics

Dependent Variable: Score

Durian	Sugar	Panelist	Mean	Std. Deviation	N
30	100	1	6.00	.	1
		2	6.00	.	1
		3	6.00	.	1
		4	6.00	.	1
		5	6.00	.	1
		6	6.00	.	1
		7	6.00	.	1
		8	7.00	.	1
		9	7.00	.	1
		10	7.00	.	1

Rows 1 through 10 of 279 (some rows are hidden)

Tests of Between-Subjects Effects

Dependent Variable: Score

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	40.800 ^a	32	1.275	3.853	.000
Intercept	4775.408	1	4775.408	14429.888	.000
Durian	16.875	1	16.875	50.991	.000
Sugar	1.408	1	1.408	4.256	.042
Block	9.842	29	.339	1.025	.447
Durian * Sugar	12.675	1	12.675	38.300	.000
Error	28.792	87	.331		
Total	4845.000	120			
Corrected Total	69.592	119			

a. R Squared = .586 (Adjusted R Squared = .434)

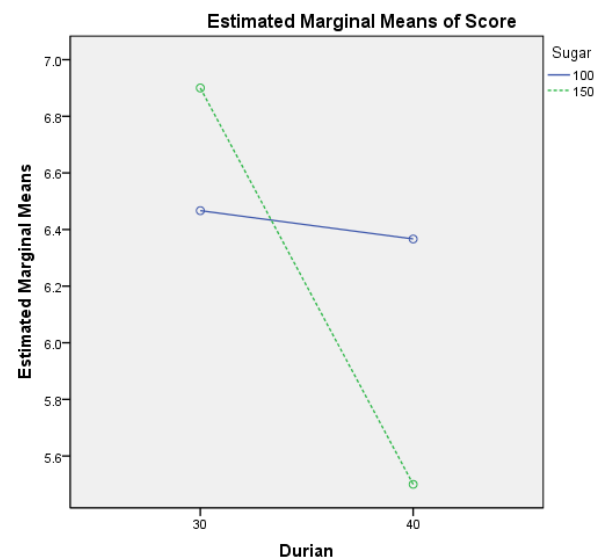
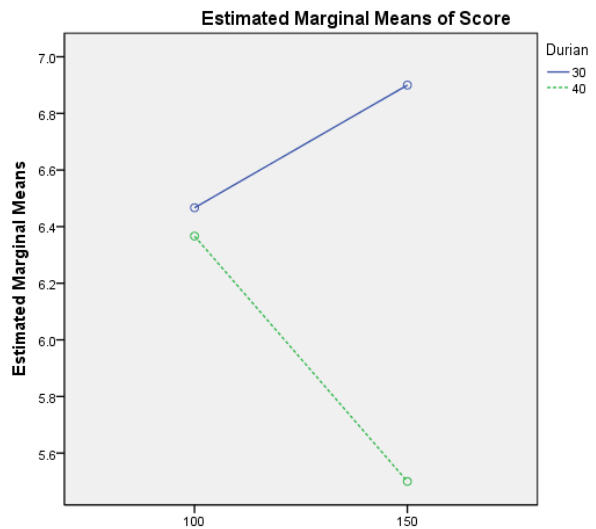
Estimated Marginal Means

Durian * Sugar

Dependent Variable: Score

Durian	Sugar	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
30	100	6.467	.105	6.258	6.675
	150	6.900	.105	6.691	7.109
40	100	6.367	.105	6.158	6.575
	150	5.500	.105	5.291	5.709

Profile Plots



ภาพที่ 10.66 ผลลัพธ์การวิเคราะห์ข้อมูลจากการทดลองแบบ 2² factorial in RCBD

ข้อสังเกต ในตาราง Between-Subjects Factors และตาราง Descriptive Statistics จะพบว่าในตารางแสดงข้อมูลไม่ครบ ให้นักวิจัยสังเกตที่ลูกศรใต้ตารางและคำอธิบายใต้ตารางที่บ่งชี้ว่ายังมีข้อมูลที่ยังซ่อนอยู่ไม่ได้แสดง ดังนั้นหากนักวิจัยต้องการเรียกดูข้อมูลทั้งหมดทำได้โดยกดคลิกที่ตารางนั้นจะปรากฏลูกศรสีแดงข้างตารางนั้น แล้วให้กดเมาส์คลิกขวาจะปรากฏกล่องคำสั่งให้เลือกคำสั่ง Set Rows to Display ... แล้วป้อนจำนวนข้อมูลที่ต้องการให้แสดง

10.7.4 การสรุปผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

ถ้านักวิจัยกำหนดระดับนัยสำคัญ (α) เท่ากับ 0.05 นักวิจัยสามารถสรุปผลลัพธ์ตามสมมติฐานที่ได้กำหนดไว้ ดังนี้

สมมติฐานข้อที่ 1 ทดสอบอิทธิพลของระดับปริมาณเนื้อทุเรียน (Main effect A)

H_0 : เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30 และ 40 กรัม มีค่าเฉลี่ยคะแนนความชอบไม่แตกต่างกัน

H_1 : เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30 และ 40 กรัม มีค่าเฉลี่ยคะแนนความชอบแตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 10.66 ค่า Sig. ของตัวแปร Durian มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ เค้กทุเรียนที่ใช้ปริมาณเนื้อทุเรียน 30 และ 40 กรัม มีค่าเฉลี่ยคะแนนความชอบแตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 2 ทดสอบอิทธิพลของระดับปริมาณน้ำตาล (Main effect B)

H_0 : เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยคะแนนความชอบไม่แตกต่างกัน

H_1 : เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยคะแนนความชอบแตกต่างกัน

จากผลลัพธ์ที่ได้ในภาพที่ 10.66 ค่า Sig. ของตัวแปร Sugar มีค่าเท่ากับ 0.042 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ เค้กทุเรียนที่ใช้ปริมาณน้ำตาล 100 และ 150 กรัม มีค่าเฉลี่ยคะแนนความชอบแตกต่างกันที่ระดับนัยสำคัญ 0.05

สมมติฐานข้อที่ 3 ทดสอบอิทธิพลร่วมระหว่างระดับปริมาณเนื้อทุเรียนและระดับปริมาณน้ำตาล (Interaction effects AxB)

H_0 : ไม่มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยคะแนนความชอบ

H_1 : มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยคะแนนความชอบ

จากผลลัพธ์ที่ได้ในภาพที่ 10.66 ค่า Sig. ของตัวแปร Durian*Sugar มีค่าเท่ากับ 0.000 ซึ่งน้อยกว่าค่า 0.05 ดังนั้น ปฏิเสธสมมติฐานหลัก (H_0) นั่นคือ มีปฏิสัมพันธ์ระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาลต่อค่าเฉลี่ยคะแนนความชอบที่ระดับนัยสำคัญ 0.05 และเมื่อพิจารณากราฟ Profile Plots พบว่า ในทุกระดับของปริมาณน้ำตาลการใช้เนื้อทุเรียน 30 กรัมมีค่าเฉลี่ยคะแนนความชอบมากกว่าเนื้อทุเรียน 40 กรัม และเมื่อพิจารณาเค้กที่ใช้เนื้อทุเรียน 30 กรัมการใช้น้ำตาล 150 กรัมมีค่าเฉลี่ยคะแนนความชอบมากกว่าที่ 100 กรัม

สมมติฐานข้อที่ 4 การทดสอบสมมติฐานอิทธิพลของบล็อก (Block)

H_0 : บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบ

H_1 : บล็อกมีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบ

จากผลลัพธ์ที่ได้ในภาพที่ 10.66 ค่า Sig. ของตัวแปร Block มีค่าเท่ากับ 0.447 ซึ่งมากกว่าค่า 0.05 ดังนั้น ยอมรับสมมติฐานหลัก (H_0) นั่นคือ บล็อกไม่มีอิทธิพลต่อค่าเฉลี่ยคะแนนความชอบที่ระดับนัยสำคัญ 0.05

10.7.5 การนำเสนอผลลัพธ์การวิเคราะห์ข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD

การเขียนรายงานผลลัพธ์การวิเคราะห์ความแปรปรวนของข้อมูลจากการวางแผนการทดลองแบบ 2^2 Factorial in RCBD ทำได้เช่นเดียวกับการวางแผนการทดลองแบบ 2^2 Factorial in CRD ในหัวข้อที่ 10.4.4.6 ซึ่งสามารถนำเสนอได้หลายแบบ ในที่นี้จะนำเสนอในรูปแบบตาราง 2 รูปแบบดังตารางที่ 10.16 และ 10.17 ซึ่งนักวิจัยจะเลือกนำเสนอแบบใดนั้นขึ้นอยู่กับวัตถุประสงค์ที่ต้องการนำเสนอข้อมูล

ตารางที่ 10.16 ค่า P-value แสดงอิทธิพลของปริมาณเนื้อทุเรียน อิทธิพลของปริมาณน้ำตาล อิทธิพลร่วมระหว่างปริมาณเนื้อทุเรียนและปริมาณน้ำตาล และอิทธิพลของบล็อก ต่อค่าเฉลี่ยคะแนนความชอบเค้กทุเรียน

อิทธิพลของปัจจัย	ค่า P-value
ปริมาณเนื้อทุเรียน	0.000*
ปริมาณน้ำตาล	0.042*
ปริมาณเนื้อทุเรียน × ปริมาณน้ำตาล	0.000*
บล็อก	0.447

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.17 ค่าเฉลี่ยคะแนนความชอบเค้กทุเรียนที่ใช้ปริมาณทุเรียนและน้ำตาลต่างกัน เมื่อใช้ผู้ทดสอบ 30 คน ประเมินด้วยสเกล 9-point hedonic scale (1 = ไม่ชอบมากที่สุด และ 9 = ชอบมากที่สุด)

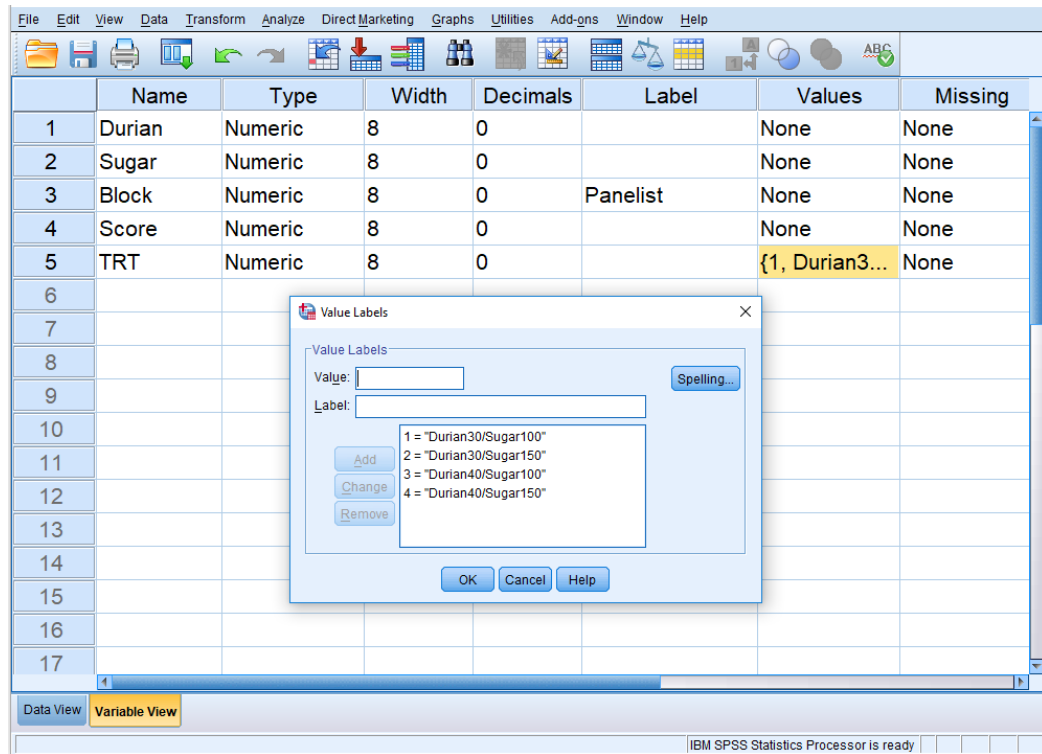
ปริมาณทุเรียน (กรัม)	ปริมาณน้ำตาล (กรัม)	คะแนนความชอบ
30	100	6.5 ± 0.5b
30	100	6.9 ± 0.8a
40	150	6.4 ± 0.5b
40	150	5.5 ± 0.5c
อิทธิพลของปัจจัย		ค่า P-value
ปริมาณเนื้อทุเรียน		0.000*
ปริมาณน้ำตาล		0.042*
ปริมาณเนื้อทุเรียน × ปริมาณน้ำตาล		0.000*
บล็อก		0.447

^{a-c} ตัวอักษรที่แตกต่างกันในแนวตั้งแสดงถึงค่าเฉลี่ยที่มีความแตกต่างกันทางสถิติ ($p \leq 0.05$)

* ระดับนัยสำคัญทางสถิติ $\alpha = 0.05$

ตารางที่ 10.17 ได้มาจากการกำหนดตัวแปร var ใหม่ชื่อ TRT ซึ่งเป็นตัวแปรที่แสดงถึงสิ่งทดลองที่เกิดจากปัจจัยร่วมระหว่างปริมาณเนื้อทุเรียน (2 ระดับ) และ ปริมาณน้ำตาล (2 ระดับ) หรือเรียกว่า “Treatment Combination” ในตัวอย่างนี้จะได้จำนวนสิ่งทดลอง (TRT) เท่ากับ 4 สิ่งทดลอง (คำนวณจากผลคูณของระดับแต่ละปัจจัย) นำข้อมูลตัวแปร TRT, ตัวแปร Block และตัวแปร Score ไปวิเคราะห์ด้วยคำสั่ง **Post Hoc...** เพื่อเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ ทำเช่นเดียวกับการวิเคราะห์ความแปรปรวนของข้อมูลที่ได้จากการวางแผนการทดลองแบบ RCBD (หัวข้อที่ 9.6.3) ตามขั้นตอนดังนี้

(1) กำหนดรายละเอียดและค่า Values ของตัวแปร TRT ในหน้าต่าง Variable View ดังภาพที่ 10.67 และป้อนข้อมูล TRT ในหน้าต่าง Data View ดังภาพที่ 10.68

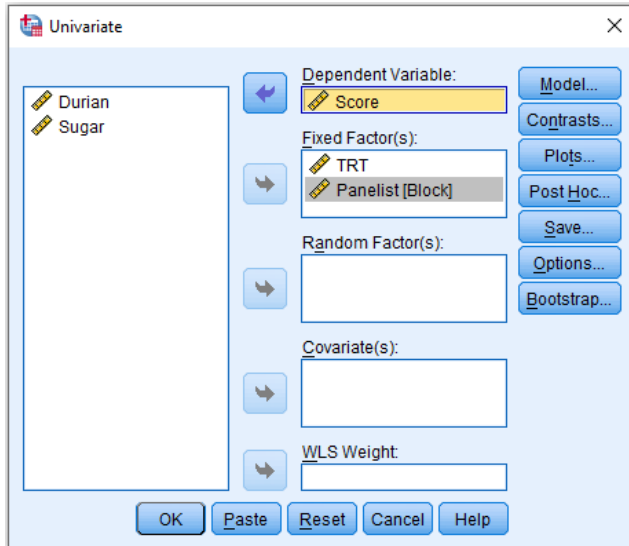


ภาพที่ 10.67 การกำหนดรายละเอียดตัวแปรใหม่ในหน้าต่าง Variable View

	Durian	Sugar	Block	Score	TRT	var	var	var	var
1	30	100	1	6	1				
2	30	100	2	6	1				
3	30	100	3	6	1				
4	30	100	4	6	1				
5	30	100	5	6	1				
6	30	100	6	6	1				
7	30	100	7	6	1				
8	30	100	8	7	1				
9	30	100	9	7	1				
10	30	100	10	7	1				
11	30	100	11	7	1				
12	30	100	12	7	1				
13	30	100	13	7	1				
14	30	100	14	7	1				
15	30	100	15	7	1				
16	30	100	16	6	1				

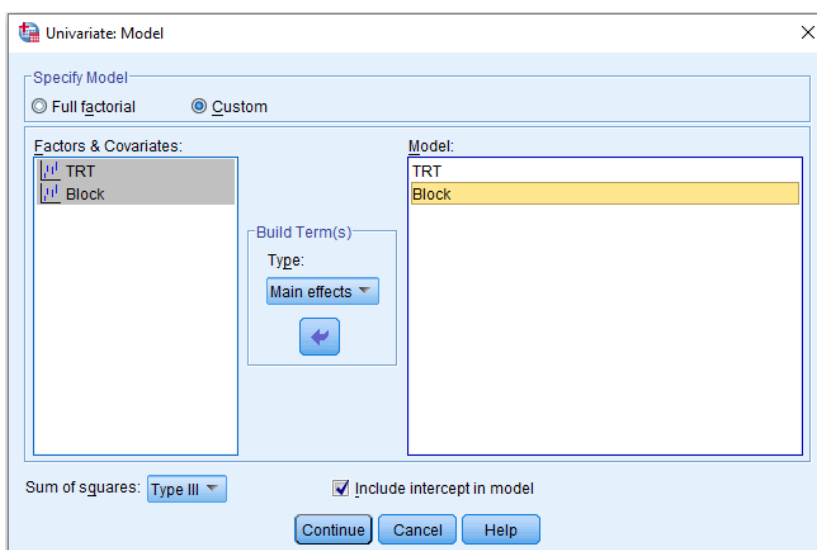
ภาพที่ 10.68 การป้อนข้อมูล TRT ในหน้าต่าง Data View

(2) เลือกเมนู **Analyze > General Linear Model > Univariate...** และเลือก ตัวแปร Score (ตัวแปรตาม) ใส่ในช่อง **Dependent Variable** และ ตัวแปร TRT และ ตัวแปร Panelist [Block] ใส่ในช่อง **Fixed Factor(s)** จะได้ดังภาพที่ 10.69



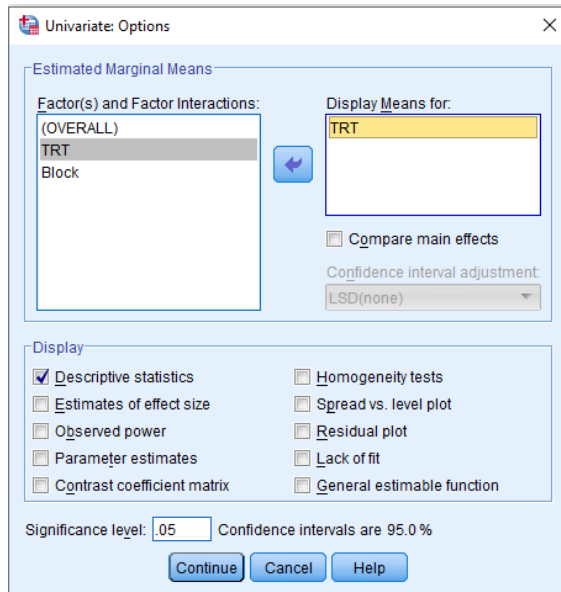
ภาพที่ 10.69 การกำหนดตัวแปรในเมนู คำสั่ง Univariate

(3) เลือกคำสั่ง **Model...** ในการวางแผนการทดลองแบบ Factorial in RCBD ต้องกำหนดตัวแบบ (Model) ชนิด **Custom** จากนั้นให้กดเลือก **Type** ของ **Build Terms(s)** เป็น **Main effects** แล้วกดเลือก ตัวแปร TRT และ ตัวแปร Block จากกล่องซ้ายมือ (Factors & Covariates) ไปใส่ในกล่องขวามือ (Model) โดยคลิกเลือกทีละตัวแปร (ดังภาพที่ 10.70) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate (ข้อสังเกต จะไม่กำหนด Build Term(s) แบบ Interaction เนื่องจากเงื่อนไขของแผนการทดลองแบบ RCBD ต้องไม่มีปฏิสัมพันธ์ระหว่างสิ่งทดลองและบล็อก)



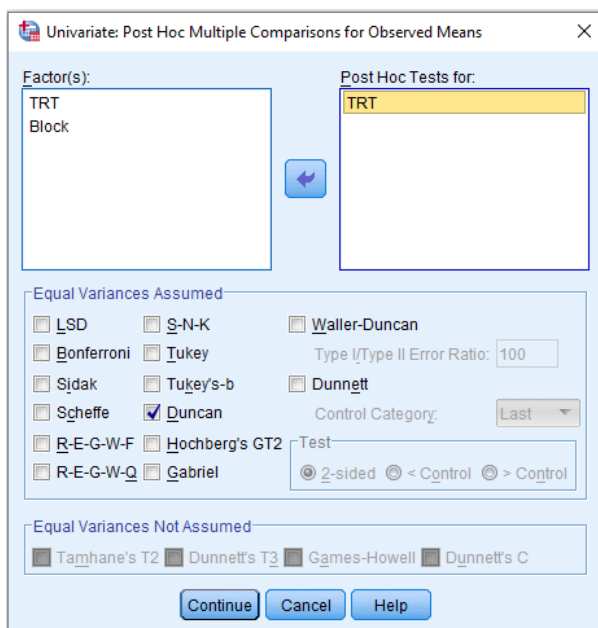
ภาพที่ 10.70 การเลือกตัวแปร ในกล่อง Model

(4) เลือกคำสั่ง **Options...** จะได้หน้าจอ ดังภาพที่ 10.71 ให้เลือก **ตัวแปร TRT** จากกล่องซ้ายมือ (Factor(s) and Factor Interactions) ไปใส่ในกล่องขวามือ (Display Means for) แล้วกดเลือก Descriptive statistics (เพื่อคำนวณค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของข้อมูลแต่ละสิ่งทดลอง) ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate



ภาพที่ 10.71 การกำหนดตัวแปรในเมนูคำสั่ง Univariate : Options

(5) เลือกคำสั่ง **Post Hoc...** เพื่อเลือกวิธีการเปรียบเทียบค่าเฉลี่ยแบบจับคู่พหุคูณ หรือเปรียบเทียบความแตกต่างของค่าเฉลี่ยแต่ละคู่ จะได้หน้าจอ ดังภาพที่ 10.72 ให้กดเลือก **ตัวแปร TRT** ใส่ในช่อง Post Hoc Tests for: และกดเลือก **วิธี Duncan** ต่อจากนั้นกด Continue เพื่อกลับมายังหน้าจอหลักของคำสั่ง Univariate และกด OK จะได้ผลลัพธ์ดังภาพที่ 10.73



ภาพที่ 10.72 การกำหนดตัวแปรและวิธีการทดสอบในเมนูคำสั่ง Post Hoc...

Univariate Analysis of Variance

Between-Subjects Factors		
	Value Label	N
TRT	1 Durian 30g, Sugar 100g	30
	2 Durian 30g, Sugar 150g	30
	3 Durian 40g, Sugar 100g	30
	4 Durian 40g, Sugar 150g	30
Panelist	1	4
	2	4
	3	4
	4	4
	5	4
	6	4
	7	4
	8	4
	9	4
	10	4
	11	4
	12	4
	13	4
	14	4
	15	4
	16	4
	17	4
	18	4
	19	4
	20	4
	21	4
	22	4
	23	4
	24	4
	25	4
	26	4
	27	4
	28	4
	29	4
	30	4

Tests of Between-Subjects Effects					
Dependent Variable: Score					
Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	40.800 ^a	32	1.275	3.853	.000
Intercept	4775.408	1	4775.408	14429.888	.000
TRT	30.958	3	10.319	31.182	.000
Block	9.842	29	.339	1.025	.447
Error	28.792	87	.331		
Total	4845.000	120			
Corrected Total	69.592	119			

a. R Squared = .586 (Adjusted R Squared = .434)

Post Hoc Tests

TRT

Homogeneous Subsets

Score				
Duncan ^{a, b}				
TRT	N	Subset		
		1	2	3
Durian 40g, Sugar 150g	30	5.50		
Durian 40g, Sugar 100g	30		6.37	
Durian 30g, Sugar 100g	30		6.47	
Durian 30g, Sugar 150g	30			6.90
Sig.		1.000	.503	1.000

Means for groups in homogeneous subsets are displayed.
Based on observed means.
The error term is Mean Square(Error) = .331.

a. Uses Harmonic Mean Sample Size = 30.000.

b. Alpha = .05.

ภาพที่ 10.73 ผลลัพธ์การวิเคราะห์ความแปรปรวนตัวแปร TRT และ ตัวแปร Score

บรรณานุกรม

- กัลยา วานิชย์บัญชา. 2548. สถิติสำหรับงานวิจัย. พิมพ์ครั้งที่ 1. กรุงเทพมหานคร : โรงพิมพ์แห่งจุฬาลงกรณ์มหาวิทยาลัย.
- กัลยา วานิชย์บัญชา. 2560. หลักสถิติ. พิมพ์ครั้งที่ 15. กรุงเทพฯ : โรงพิมพ์สามลดา.
- กัลยา วานิชย์บัญชา. 2560. การใช้ SPSS for Windows ในการวิเคราะห์ข้อมูล. พิมพ์ครั้งที่ 30. กรุงเทพฯ : โรงพิมพ์สามลดา.
- คณาจารย์ภาควิชาพัฒนาผลิตภัณฑ์ มหาวิทยาลัยเกษตรศาสตร์. 2559. การพัฒนาผลิตภัณฑ์ในอุตสาหกรรมเกษตร. กรุงเทพฯ : สำนักพิมพ์มหาวิทยาลัยเกษตรศาสตร์.
- ฉัตยาพร เสมอใจ. พฤติกรรมผู้บริโภค. กรุงเทพฯ : โรงพิมพ์วีพรีน(1991).
- ชัยวิชิต เขียวชนะ. 2560. สถิติสำหรับการวิจัย : แนวคิดและการประยุกต์ใช้. พิมพ์ครั้งที่ 1. กรุงเทพฯ : โรงพิมพ์แห่งจุฬาลงกรณ์มหาวิทยาลัย.
- ปิติพร ฤทธิเรืองเดช. 2558. เอกสารประกอบการสอนวิชา 01054351 : หลักการพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร. กรุงเทพฯ : ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเกษตรศาสตร์.
- เพ็ญขวัญ ชมปรีดา. 2556. การประเมินคุณภาพทางประสาทสัมผัสและการยอมรับของผู้บริโภค. พิมพ์ครั้งที่ 2. โรงพิมพ์วิสตา อินเตอร์พรีนซ์.
- ยุทธ ไถยวรรณ. 2559. การวางแผนการทดลองสำหรับการวิจัย. กรุงเทพฯ : สำนักพิมพ์แห่งจุฬาลงกรณ์มหาวิทยาลัย.
- วิชัย หฤทัยธนาสันต์. 2527. เอกสารประกอบคำสอนวิชา PD351 การพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร. กรุงเทพฯ : ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเกษตรศาสตร์.
- ศิริชัย พงษ์วิชัย. 2558. การวิเคราะห์ข้อมูลทางสถิติด้วยคอมพิวเตอร์. พิมพ์ครั้งที่ 25. กรุงเทพฯ : โรงพิมพ์แห่งจุฬาลงกรณ์มหาวิทยาลัย.
- ศิริวรรณ เสรีรัตน์, สมชาย หิรัญกิตติ, วลัยลักษณ์ อัครธีรวงศ์, ชวลิต ประภาวนนท์, จิรศักดิ์ จิยะจันทร์ และ ณา จันทรสม. 2540. การวิจัยการตลาด. กรุงเทพฯ : โรงพิมพ์เอ.เอ็น.การพิมพ์.
- ศิริวรรณ เสรีรัตน์, ศุภร เสรีรัตน์, องอาจ ปทะวานิช, ปริญ ลักษิตานนท์ และ สุพีร์ ลิ้มไทย. 2543. หลักการตลาด. กรุงเทพฯ : โรงพิมพ์ธีระฟิล์มและไซเท็กซ์.
- สายชล สีนสมบูรณ์ทอง. 2559. การวิเคราะห์ตัวแปรหลายตัวโดยใช้ SPSS และ MINITAB. พิมพ์ครั้งที่ 1. กรุงเทพฯ : จามจุรีโปรดักส์.

- สายชล สันสมบูรณ์ทอง. 2558. สถิติเบื้องต้น. พิมพ์ครั้งที่ 11. กรุงเทพฯ : จามจุรีโปรดักส์.
- สายชล สันสมบูรณ์ทอง. 2558. การวางแผนแบบการทดลอง เล่ม 1. พิมพ์ครั้งที่ 1. กรุงเทพฯ : จามจุรีโปรดักส์.
- เสรี วงษ์มณฑา. 2542. พฤติกรรมผู้บริโภค. กรุงเทพฯ : โรงพิมพ์ธีระฟิล์มและไซเท็กซ์.
- อนุวัตร แจ่มชัด. 2558. สถิติสำหรับการพัฒนาผลิตภัณฑ์และการประยุกต์. กรุงเทพฯ : ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเกษตรศาสตร์.
- อิสรพงษ์ พงษ์ศิริกุล. 2545. การวิเคราะห์ผลทางสถิติโดยใช้โปรแกรมสำเร็จรูปสำหรับอุตสาหกรรมเกษตร. พิมพ์ครั้งที่ 3. เชียงใหม่ : ภาควิชาเทคโนโลยีการพัฒนาระบบผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเชียงใหม่.
- Capenter, R.P., Lyon, D.H., and Hasdell, T.A. 2000. Guidelines for sensory analysis in food product development and quality control. 2nd ed. Maryland : Aspen Publishers, Inc.
- Earle M.D., Earle, R.L., and Anderson, A.A. 2001. Product development. UK : Woodhead Publishing Limited.
- Gravetter, F.J., and Wallnau, L.B. 2007. Statistics for the behavioral sciences. 7th ed. Australia : Thomson Wadsworth.
- Healey, J.F. 2009. Statistics : A tool for social research. 8th ed. Australia : Wadsworth.
- Montgomery, D.C. 2013. Design and analysis of experiments. 8th ed. New York : John Wiley and Sons.
- Meilgaard, M., Civille, G.V., and Carr, B.T. 1999. Sensory Evaluation Techniques. 3rd ed. New York : CRC Press.
- Rodrigues, M.I., and Antonio, F.I. 2015. Experimental design and process optimization. London : CRC Press.
- Sprinthall, R.C. 2003. Basic statistical analysis. Boston: Allyn and Bacon.
- Yam, K.L., Takhistov, P.T., and Miltz, J. 2005. Intelligent Packaging : Concepts and Applications. J. Food Sci. 70 : R1-R10.

ดรรชนี

ก		การวิเคราะห์ความแปรปรวนแบบ	182, 208
การจัดกลุ่มสิ่งทดลอง	46, 302	สองทาง	
การกำหนดค่าของตัวแปร	57	การวิเคราะห์ความแปรปรวนแบบ	182
การจัดการวงจรชีวิตผลิตภัณฑ์	27	หลายทาง	
การทดลองขั้นวิภูติ	42	การวิเคราะห์สถานการณ์	18
การทดลองขั้นสาธิต	42	การสร้างแนวคิดผลิตภัณฑ์	20
การทดลองเบื้องต้น	42	การสุ่ม	49
การทดลองแบบแฟกทอเรียล	301		
การทดสอบที่มีนัยสำคัญ	118	ข	
การทดสอบที่ไม่มีนัยสำคัญ	118	ข้อกำหนดผลิตภัณฑ์	23
การทดสอบสมมติฐาน	118	ข้อมูลเชิงคุณภาพ	37
การทดสอบสมมติฐานแบบทางเดียว	121	ข้อมูลเชิงปริมาณ	37
การทดสอบสมมติฐานแบบสองทาง	119	ข้อมูลทฤษฎี	36
การทดสอบสมมติฐานแบบสองทาง	119	ข้อมูลปฐมภูมิ	36
การทดสอบสมมติฐานแบบทางเดียว	121	เขตวิฤต	124
การทำซ้ำ	48		
การบล็อก	50, 272	ค	
การบันทึกผลลัพธ์	69	ความคลาดเคลื่อนในการทดลอง	48
การป้อนข้อมูล	60	ความแปรปรวนที่อธิบายได้	180, 181
การแปรรูปขั้นต่ำ	6	ความแปรปรวนที่อธิบายไม่ได้	180, 181
การวางแผนการทดลองแบบสุ่มใน	271	ค่ากลางของข้อมูล	93
บล็อกสมบูรณ์		ค่าเฉลี่ย	93
การวางแผนการทดลองแบบสุ่ม	243	ค่าฐานนิยม	93
สมบูรณ์		ค่าแปรปรวน	93
การวิเคราะห์ความแปรปรวนแบบ	182, 186	ค่าพิสัย	93
ทางเดียว		ค่ามัธยฐาน	93
การวิเคราะห์ความถดถอยเชิงพหุ	134	ค่าระดับนัยสำคัญ	122
การวิเคราะห์ความแปรปรวน	179	ความคลาดเคลื่อนในการทดลอง	48

ความแปรปรวนที่อธิบายได้	180, 181	ผ	
ค่าวิกฤต	124	ผลตอบแทน	43
		ผลิตภัณฑ์คงทน	10
ด		ผลิตภัณฑ์ต้นแบบ	25
ตัวแปรเชิงคุณภาพ	35	ผลิตภัณฑ์ที่จับต้องได้	10
ตัวแปรเชิงปริมาณ	35	ผลิตภัณฑ์ที่จับต้องไม่ได้	10
ตัวแปรต้น	35	ผลิตภัณฑ์ที่ไม่ได้แสวงซื้อ	11
ตัวแปรต่อเนื่อง	35	ผลิตภัณฑ์ในรูปของบริการ	10
ตัวแปรตาม	35	ผลิตภัณฑ์เพื่อผู้บริโภค	10
ตัวแปรไม่ต่อเนื่อง	35	ผลิตภัณฑ์เพื่ออุตสาหกรรม	11
ตัวแปรอิสระ	35	ผลิตภัณฑ์ไม่คงทน	10
ตัวอย่าง	34	ผลิตภัณฑ์สะดวกซื้อ	11
ตัวอย่างที่เป็นอิสระกัน	145	ผลิตภัณฑ์ใหม่	13
ตัวอย่างที่ไม่เป็นอิสระกัน	146	ผลิตภัณฑ์เจาะจงซื้อ	11
ตารางการแจกแจงความถี่ร่วม	102	ผลิตภัณฑ์เลือกซื้อ	11
		แผนภูมิแก๊งปลา	302
ท		ร	
ทรีทเมนต์ควบคุม	43	ระดับของปัจจัย	46
บ		ระดับของผลิตภัณฑ์	12
บรรจุภัณฑ์ฉลาด	6	รูปแบบของคำถามใน	73
บรรจุภัณฑ์แบบแอคทีฟ	6	แบบสอบถาม	
บรรจุภัณฑ์อินเทลลิเจนท์	6		
ป		ว	
ปฏิสัมพันธ์ระหว่างปัจจัย	308	วงจรชีวิตผลิตภัณฑ์	3
ประชากร	34	วิจัยอุตสาหกรรมประยุกต์	2
ประเภทของผลิตภัณฑ์ใหม่	13	วิธีการทดสอบพหุคูณ	184
ประเภทและระดับของผลิตภัณฑ์	10		
ปัจจัย	46	ส	
ปัจจัยหลัก	208	สเกลนามกำหนด	37
		สเกลนามบัญญัติ	37

สเกลแบบช่วง	39	อ	
สเกลอัตราส่วน	39	อิทธิพลร่วมระหว่างปัจจัย	308
สเกลอันดับ	38	อิทธิพลหลัก	308
สเกลอันตรภาค	39		
สถิติเชิงพรรณนา	72		
สถิติเชิงพรรณนา	32		
สถิติเชิงพรรณนาสำหรับข้อมูลเชิง	72, 78		
คุณภาพหรือเชิงกลุ่ม			
สถิติเชิงพรรณนาสำหรับข้อมูลเชิง	73, 93		
ปริมาณ			
สถิติทดสอบ	33		
สถิติประมาณ	33		
สถิติอนุमान	33, 116		
สมมติฐานการวิจัย	117		
สมมติฐานทางเลือก	117		
สมมติฐานทางสถิติ	117		
สมมติฐานรอง	117		
สมมติฐานรองแบบมีทิศทาง	118		
สมมติฐานรองแบบไม่มีทิศทาง	118		
สมมติฐานว่าง	117		
สมมติฐานวิจัย	116		
สมมติฐานหลัก	117		
สิ่งทดลอง	43		
สิ่งทดลองควบคุม	43		
หน่วยทดลองกลุ่ม	45		
หน่วยทดลองเดี่ยว	44		
แหล่งของความแปรปรวน	180		

ประวัติผู้เขียน



รองศาสตราจารย์ ดร.ปิติพร ฤทธิเรืองเดช

สถานที่ทำงาน

ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร มหาวิทยาลัยเกษตรศาสตร์

ประวัติการศึกษา

- วท.บ. (พัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร) เกียรตินิยมอันดับสอง ม. เกษตรศาสตร์
- วท.ม. (พัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร) ม. เกษตรศาสตร์
- ประ.ด. (พัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร) ม. เกษตรศาสตร์

สาขาวิชาที่มีความเชี่ยวชาญพิเศษ

- การพัฒนาผลิตภัณฑ์อาหารจากวัตถุดิบทางการเกษตร และการวิเคราะห์คุณภาพด้านเนื้อสัมผัส
- การประยุกต์ใช้สถิติในงานพัฒนาผลิตภัณฑ์ทางอุตสาหกรรมเกษตร
- การประยุกต์ใช้เทคนิคสเปกโทรสโกปีย่านใกล้อินฟราเรด (Near-infrared spectroscopy, NIRs) วิเคราะห์คุณภาพผลผลิตทางการเกษตรและผลิตภัณฑ์อาหารและกึ่งอาหาร

ประสบการณ์การทำงาน

- อาจารย์ประจำสอนในระดับปริญญาตรีและบัณฑิตศึกษา เช่น รายวิชา อุตสาหกรรมเกษตรเบื้องต้น, เทคนิคสำหรับพัฒนาผลิตภัณฑ์, หลักการพัฒนาผลิตภัณฑ์อุตสาหกรรมเกษตร, สถิติสำหรับการพัฒนาผลิตภัณฑ์, สถิติประยุกต์สำหรับการพัฒนาผลิตภัณฑ์, ระเบียบวิธีวิจัยพื้นฐานทางพัฒนาผลิตภัณฑ์ อุตสาหกรรมเกษตร และ การพัฒนาผลิตภัณฑ์จากเมล็ดพืชและพืชหัว
- นักวิชาการผู้ตรวจประเมินห้องปฏิบัติการตามมาตรฐาน ISO/IEC 17025: 2005 ของสำนักมาตรฐานห้องปฏิบัติการ กรมวิทยาศาสตร์การแพทย์ กระทรวงสาธารณสุข
- อาจารย์พิเศษและวิทยากรบรรยายในหน่วยงานและมหาวิทยาลัยต่างๆ ทั้งภาครัฐและเอกชน ในหัวข้อเรื่องต่างๆ เช่น การพัฒนาผลิตภัณฑ์ทางอุตสาหกรรมเกษตร, การประเมินคุณภาพทางประสาทสัมผัส และ การใช้เทคนิคสเปกโทรสโกปีย่านใกล้อินฟราเรดในการตรวจสอบคุณภาพผลผลิตทางการเกษตรและผลิตภัณฑ์อุตสาหกรรมเกษตร

หนังสือ “การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรม SPSS” เล่มนี้ จัดพิมพ์เป็นครั้งที่ 3 มีการปรับปรุงแก้ไขจากการพิมพ์ในครั้งก่อนให้มีความสมบูรณ์มากยิ่งขึ้น โดยยังคงเนื้อหาสำคัญภายในเล่มซึ่งได้อธิบายเทคนิคทางสถิติที่มีความสำคัญต่อการพัฒนาผลิตภัณฑ์ ในอุตสาหกรรมเกษตร ประกอบด้วย กฎเกณฑ์ทางสถิติและการประยุกต์ใช้ในงานพัฒนาผลิตภัณฑ์ ตลอดจนการใช้โปรแกรมสำเร็จรูปทางสถิติ SPSS ในการวิเคราะห์ข้อมูลงานวิจัยและพัฒนา เริ่มตั้งแต่ การป้อนข้อมูล การเลือกใช้ค่าสิ่งสถิติเพื่อวิเคราะห์ การแปลผล และการเขียนรายงานผลการวิเคราะห์ทางสถิติ หนังสือเล่มนี้เหมาะสำหรับนิสิตนักศึกษาใช้ประกอบการเรียนการสอนในระดับมหาวิทยาลัยและบุคลากรที่ทำงานในด้านการวิจัยและพัฒนาผลิตภัณฑ์ทางด้านอุตสาหกรรมเกษตร

ISBN 978-616-586-105-2



ราคา 350 บาท



การวิเคราะห์ข้อมูลงานวิจัยและพัฒนาผลิตภัณฑ์โดยใช้โปรแกรม SPSS รศ.ดร.ปิณฑร ฤทธิเรืองเดช

การวิเคราะห์ข้อมูลงานวิจัยและ พัฒนาผลิตภัณฑ์โดยใช้โปรแกรม

SPSS

Data Analysis for Research and Product Development using SPSS

รองศาสตราจารย์ ดร.ปิณฑร ฤทธิเรืองเดช
ภาควิชาพัฒนาผลิตภัณฑ์ คณะอุตสาหกรรมเกษตร
มหาวิทยาลัยเกษตรศาสตร์