

การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่อง สำหรับการจัดการซัพพลายเชน

DATA ANALYTICS AND MACHINE LEARNING FOR
SUPPLY CHAIN MANAGEMENT



ผศ.ดร.ณัฐพัชร อารีรัชกุลกานต์

วิทยาลัยโลจิสติกส์และซัพพลายเชน
มหาวิทยาลัยราชภัฏสวนสุนันทา
2568



การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องสำหรับ
การจัดการซัพพลายเชน

DATA ANALYTICS AND MACHINE LEARNING FOR
SUPPLY CHAIN MANAGEMENT

ผศ.ดร.ณัฐพัชร อารีรัชกุลกานต์

วิทยาลัยโลจิสติกส์และซัพพลายเชน
มหาวิทยาลัยราชภัฏสวนสุนันทา

2568

ชื่อหนังสือ การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องสำหรับการจัดการซัพพลายเชน
(Data Analytics and machine learning for supply chain management)

ผู้เรียบเรียง ผศ.ดร. ณัฐพัชร์ อารีรัชกุลกานต์

ข้อมูลทางบรรณานุกรมของหอสมุดแห่งชาติ

National Library of Thailand Cataloging in Publication Data

ณัฐพัชร์ อารีรัชกุลกานต์.

การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องสำหรับการจัดการซัพพลายเชน = Data Analytics and machine learning for supply chain management.-- นครปฐม : สาขาวิชาการจัดการโลจิสติกส์ และซัพพลายเชน วิทยาลัยโลจิสติกส์และซัพพลายเชน มหาวิทยาลัยราชภัฏสวนสุนันทา, 2569.
341 หน้า.

1.การเรียนรู้ของเครื่อง. 2. การจัดการห่วงโซ่อุปทาน. I. ชื่อเรื่อง.

658.5

ISBN (e-book) 978-616-636-629-7

สงวนลิขสิทธิ์ตามพระราชบัญญัติลิขสิทธิ์

โดย ผู้ช่วยศาสตราจารย์ ดร.ณัฐพัชร์ อารีรัชกุลกานต์ © พ.ศ. 2568 ห้ามคัดลอก ลอกเลียน ทำซ้ำ ดัดแปลง จัดพิมพ์
เผยแพร่หรือจำหน่ายส่วนใดส่วนหนึ่งของหนังสืออิเล็กทรอนิกส์นี้ เว้นแต่จะได้รับอนุญาตเป็นลายลักษณ์อักษร
จากเจ้าของลิขสิทธิ์

พิมพ์ครั้งที่ 1 2568

เผยแพร่โดย สำนักพิมพ์มหาวิทยาลัยเทคโนโลยีราชมงคลสุวรรณภูมิ (RUS Digital Publishing)
เลขที่ 60 หมู่ 3 อาคารสำนักวิทยบริการและเทคโนโลยีสารสนเทศ ถนนสายเอเชีย
(กรุงเทพฯ - นครสวรรค์) ตำบลหันตรา อำเภอพระนครศรีอยุธยา
จังหวัดพระนครศรีอยุธยา 13000

เบอร์โทร 088-9216914

E-mail digitalpub@rmutsb.ac.th

Website <https://dpub.rmutsb.ac.th>

คำนำ

ตำราเรื่อง การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องสำหรับการจัดการซัพพลายเชน (Data Analytics and Machine Learning for Supply Chain Management) จัดทำขึ้นเพื่อใช้เป็นส่วนหนึ่งสำหรับประกอบการเรียนการสอนในรายวิชา BUA6104 วิธีการวิเคราะห์ระบบการจัดการโลจิสติกส์และซัพพลายเชน (Analytic Methods of Logistics and Supply Chain Management System) ในภาคเรียนที่ 3/2568 โดยผู้เขียนได้เรียบเรียงเนื้อหาอย่างเป็นระบบ ครอบคลุมเนื้อหาสาระโดยเน้นในส่วนของการประยุกต์การเรียนรู้ของเครื่องในการวิเคราะห์ซัพพลายเชนและโลจิสติกส์ในการทำงานจริง สำหรับนักศึกษาระดับปริญญาโท ของหลักสูตรบริหารธุรกิจมหาบัณฑิต สาขาวิชาการจัดการโลจิสติกส์และซัพพลายเชน วิทยาลัยโลจิสติกส์และซัพพลายเชน มหาวิทยาลัยราชภัฏสวนสุนันทา เพื่อใช้เป็นเครื่องมือสำคัญที่มุ่งเน้นให้ผู้เรียนมีความรู้ความเข้าใจในเนื้อหา เนื้อหาภายในตำราแบ่งออกเป็น 5 ส่วน ได้แก่

ส่วนที่ 1 พื้นฐานการวิเคราะห์ข้อมูลซัพพลายเชน

ส่วนที่ 2 การวิเคราะห์เชิงพรรณนาและเชิงวินิจัย

ส่วนที่ 3 การวิเคราะห์เชิงพยากรณ์

ส่วนที่ 4 การวิเคราะห์เพื่อการตัดสินใจ

และ ส่วนที่ 5 กรณีศึกษา/งานวิจัยที่เกี่ยวข้องและน่าสนใจ

รูปภาพและตารางในตำราเล่มนี้ ในส่วนที่ไม่ได้มีการระบุแหล่งที่มาจะจัดทำขึ้นโดยผู้เขียนเอง โดยใช้โปรแกรม Microsoft Word และ Microsoft Excel เป็นหลัก สำหรับการพัฒนาโปรแกรมและการประมวลผลข้อมูล ใช้ภาษา Python อย่างไรก็ตาม รูปภาพและตารางบางส่วน โดยเฉพาะในส่วนที่ 5 กรณีศึกษา/งานวิจัยที่เกี่ยวข้องและน่าสนใจ ส่วนใหญ่ถูกสร้างขึ้นโดยใช้โปรแกรม Python โดยเฉพาะ เพื่อสนับสนุนการวิเคราะห์และการนำเสนอข้อมูลให้มีความชัดเจนมากยิ่งขึ้น

อย่างไรก็ดี หวังว่าตำรานี้คงอำนวยความสะดวกต่อการเรียนการสอนตามสมควร หากท่านที่นำไปใช้มีข้อเสนอแนะ ผู้เขียนยินดีรับฟังข้อคิดเห็นต่าง ๆ และขอขอบคุณมา ณ โอกาสนี้

ผศ.ดร.ณัฐพัชร์ อารีรัชกุลกานต์

สารบัญ

	หน้า
คำนำ.....	ก
สารบัญ.....	ข
ส่วนที่ 1 พื้นฐานการวิเคราะห์ข้อมูลซัพพลายเชน	
บทที่ 1 หลักการวิเคราะห์ข้อมูลซัพพลายเชน.....	1
1.1 วิทยาการข้อมูล.....	1
1.2 ความเกี่ยวข้องกันระหว่างวิทยาการข้อมูลและคลาวด์คอมพิวติ้ง.....	8
1.3 การจัดการซัพพลายเชน (Supply Chain Management, SCM).....	10
1.4 องค์ประกอบของการวิเคราะห์ข้อมูลซัพพลายเชน.....	13
1.5 หลักการวิเคราะห์ซัพพลายเชนอย่างมีประสิทธิภาพ.....	15
สรุป.....	16
แบบฝึกหัด.....	17
เอกสารอ้างอิง.....	18
บทที่ 2 ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล.....	20
2.1 ข้อมูลและคุณค่าของข้อมูลในการจัดการซัพพลายเชน.....	21
2.2 ความแตกต่างของซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล.....	24
2.3 เทคโนโลยีที่ทำให้เกิดข้อมูลขนาดใหญ่ในซัพพลายเชน.....	27
2.4 คุณลักษณะของข้อมูลขนาดใหญ่ (Big Data).....	32
2.5 ข้อมูลขนาดใหญ่ (Big Data) ในซัพพลายเชน.....	34
สรุป.....	37
แบบฝึกหัด.....	38
เอกสารอ้างอิง.....	39
ส่วนที่ 2 การวิเคราะห์เชิงพรรณนาและเชิงวินิจัย	
บทที่ 3 การวิเคราะห์ข้อมูลลูกค้า.....	42
3.1 ลูกค้าในซัพพลายเชนและการทำความเข้าใจลูกค้า.....	43
3.2 การวิเคราะห์กลุ่มลูกค้า (Cohort analysis).....	47
3.3 การวิเคราะห์ RFM.....	52
3.4 อัลกอริทึมการจัดกลุ่ม (Clustering Algorithms) เพื่อทำการแบ่งกลุ่มลูกค้า... ..	58
3.5 การกำหนดนโยบายตามการแบ่งกลุ่มลูกค้า.....	62

สารบัญ (ต่อ)

	หน้า
สรุป.....	64
แบบฝึกหัด.....	65
เอกสารอ้างอิง.....	67
บทที่ 4 การวิเคราะห์ข้อมูลจัดซื้อจัดจ้าง.....	69
4.1 การจัดซื้อจัดจ้าง.....	70
4.2 การคัดเลือกซัพพลายเออร์.....	72
4.3 การประเมินซัพพลายเออร์.....	76
4.4 การจัดการความสัมพันธ์กับซัพพลายเออร์.....	88
4.5 การบริหารความเสี่ยง.....	90
สรุป.....	94
แบบฝึกหัด.....	95
เอกสารอ้างอิง.....	97
ส่วนที่ 3 การวิเคราะห์เชิงพยากรณ์	
บทที่ 5 การจัดการและพยากรณ์ความต้องการสินค้า.....	100
5.1 การจัดการความต้องการสินค้า.....	101
5.2 การพยากรณ์ความต้องการสินค้าและรูปแบบข้อมูลอนุกรมเวลา.....	103
5.3 การพยากรณ์ด้วยการวิเคราะห์อนุกรมเวลาแบบคลาสสิก.....	107
5.4 การพยากรณ์ด้วยการวิเคราะห์การถดถอยอัตโนมัติ.....	110
5.5 การประเมินความคลาดเคลื่อนของการพยากรณ์.....	113
สรุป.....	117
แบบฝึกหัด.....	119
เอกสารอ้างอิง.....	121
บทที่ 6 การพยากรณ์ความต้องการสินค้าด้วยวิธีการเรียนรู้ของเครื่อง.....	123
6.1 แบบจำลองการถดถอยเชิงเส้น (Linear Regression: LR).....	125
6.2 แบบจำลองป่าสุ่ม (Random Forest: RF).....	133
6.3 แบบจำลองโครงข่ายประสาทเทียม (Artificial Neural Network: ANN).....	139
สรุป.....	145
แบบฝึกหัด.....	147
เอกสารอ้างอิง.....	149

สารบัญ (ต่อ)

	หน้า
ส่วนที่ 4 การวิเคราะห์เพื่อการตัดสินใจ	
บทที่ 7 การจัดการและวิเคราะห์ข้อมูลโลจิสติกส์.....	151
7.1 การจัดการโลจิสติกส์.....	152
7.2 กิจกรรมหลักในการจัดการโลจิสติกส์.....	154
7.3 รูปแบบการขนส่ง.....	156
7.4 การออกแบบเครือข่ายโลจิสติกส์.....	158
7.5 การออกแบบเครือข่ายโลจิสติกส์โดยใช้แบบจำลองเชิงเส้น.....	166
สรุป.....	172
แบบฝึกหัด.....	173
เอกสารอ้างอิง.....	176
บทที่ 8 การวิเคราะห์ข้อมูลโลจิสติกส์เพื่อจัดเส้นทางขนส่ง.....	179
8.1 คุณลักษณะของปัญหาการจัดเส้นทางขนส่ง.....	180
8.2 การจัดเส้นทางขนส่ง.....	186
8.3 การใช้โปรแกรมเชิงเส้นตรงแก้ปัญหาการจัดเส้นทางขนส่ง.....	187
8.4 การใช้วิธีฮิวริสติก (Heuristic) แก้ปัญหาการจัดเส้นทางขนส่ง.....	194
8.5 การจัดเส้นทางขนส่งด้วยการเรียนรู้ของเครื่อง.....	200
สรุป.....	207
แบบฝึกหัด.....	208
เอกสารอ้างอิง.....	210
ส่วนที่ 5 กรณีศึกษา/งานวิจัยที่เกี่ยวข้องและน่าสนใจ	
บทที่ 9 การประยุกต์ใช้การเรียนรู้ของเครื่องในการจัดกลุ่มลูกค้า.....	213
9.1 วัตถุประสงค์ของกรณีศึกษา.....	214
9.2 บริบทธุรกิจ.....	214
9.3 ชุดข้อมูลและตัวแปรที่ใช้.....	215
9.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง.....	216
9.5 สรุปผลลัพธ์จากการจัดกลุ่มลูกค้า.....	234
เอกสารอ้างอิง.....	236

สารบัญ (ต่อ)

	หน้า
บทที่ 10 กรณีศึกษาการพยากรณ์อุปสงค์ด้วยวิธีการเรียนรู้ของเครื่อง.....	237
10.1 วัตถุประสงค์ของกรณีศึกษา.....	239
10.2 บริบทธุรกิจ.....	240
10.3 ชุดข้อมูลและตัวแปรที่ใช้.....	240
10.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง.....	242
10.5 สรุปผลลัพธ์จากการพยากรณ์อุปสงค์.....	254
เอกสารอ้างอิง.....	255
บทที่ 11 จากทฤษฎีสู่การปฏิบัติ: การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อ.....	256
การจัดเส้นทางการขนส่ง	
11.1 วัตถุประสงค์ของกรณีศึกษา.....	257
11.2 บริบทธุรกิจ.....	258
11.3 ชุดข้อมูลและตัวแปรที่ใช้.....	258
11.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง.....	259
11.5 สรุปผลลัพธ์จากการจัดเส้นทางการขนส่ง.....	274
เอกสารอ้างอิง.....	275
บทที่ 12 งานวิจัยที่น่าสนใจเกี่ยวกับการคัดเลือกซัพพลายเออร์อย่างยั่งยืนด้วย.....	276
กระบวนการวิเคราะห์เชิงลำดับชั้น	
12.1 วัตถุประสงค์ของงานวิจัย.....	277
12.2 บริบทธุรกิจ.....	278
12.3 ชุดข้อมูลและตัวแปรที่ใช้.....	279
12.4 ขั้นตอนการวิเคราะห์ด้วยวิธี AHP.....	280
12.5 สรุปผลงานวิจัย.....	288
เอกสารอ้างอิง.....	288

สารบัญ (ต่อ)

	หน้า
บทที่ 13 งานวิจัยที่น่าสนใจเกี่ยวกับการพยากรณ์สินค้าใหม่ด้วยวิธีโครงข่ายประสาทเทียม	291
13.1 บริบทและปัญหาของกรณีศึกษา.....	292
13.2 การวิเคราะห์ข้อมูลและพัฒนาแบบจำลอง.....	293
13.3 การประเมินและเปรียบเทียบผลการพยากรณ์.....	297
13.4 การประยุกต์ใช้ผลการพยากรณ์ในซัพพลายเชน.....	299
13.5 บทเรียนจากงานวิจัยและสรุป.....	301
เอกสารอ้างอิง.....	303
บรรณานุกรม.....	304
ดัชนีท้ายเล่ม.....	320

สารบัญตาราง

ตารางที่	หน้า
2.1 ตารางเปรียบเทียบแนวทางการดำเนินการ.....	25
3.1 ตารางตัวอย่างการคำนวณค่า RFM.....	55
3.2 ตารางตัวอย่างการแปลงค่า RFM.....	56
3.3 ตารางตัวอย่างการแบ่งกลุ่มลูกค้าด้วยการใช้ค่า RFM.....	57
3.4 ตารางข้อมูลสำหรับตัวอย่างการจัดกลุ่มลูกค้าด้วยอัลกอริทึม K-Means.....	60
3.5 ตัวอย่างการคำนวณ.....	61
4.1 ข้อเปรียบเทียบระหว่าง Vertical Integration และ Outsourcing.....	71
4.2 สเกลการเปรียบเทียบเชิงคู่ใน AHP.....	80
4.3 เมตริกการเปรียบเทียบเชิงคู่.....	80
4.4 การแปลงคะแนนการเปรียบเทียบเป็นคะแนนมาตรฐาน.....	82
4.5 ค่า RI สำหรับวิธี AHP.....	83
4.6 ค่าคะแนนความสอดคล้อง.....	84
4.7 ค่าคะแนนการเปรียบเทียบเชิงคู่เกณฑ์การประเมิน.....	84
4.8 ค่าคะแนนประเมินชีพพลายเออร์ตามราคา.....	85
4.9 ค่าคะแนนประเมินชีพพลายเออร์ตามคุณภาพ.....	85
4.10 ค่าคะแนนประเมินชีพพลายเออร์ตามความเร็ว.....	85
4.11 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้องของเกณฑ์.....	86
4.12 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้องตามราคา.....	86
4.13 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้องตามคุณภาพ.....	87
4.14 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้องตามความเร็ว.....	87
4.15 ค่าคะแนนถ่วงน้ำหนักของชีพพลายเออร์แต่ละราย.....	87
6.1 ข้อมูลยอดขายและปัจจัยที่เกี่ยวข้องของปี ค.ศ. 2020.....	130
6.2 ค่าสัมประสิทธิ์.....	131
6.3 ข้อมูลชุดฝึก.....	136
6.4 ข้อมูลสุ่มตัวอย่างของต้นไม้ต้นที่ 1.....	137
6.5 ข้อมูลสุ่มตัวอย่างของต้นไม้ต้นที่ 2.....	137
6.6 ข้อมูลสุ่มตัวอย่างของต้นไม้ต้นที่ 3.....	137
6.7 ค่าน้ำหนักและค่าอคติเริ่มต้น.....	141
7.1 บทบาทการจัดการโลจิสติกส์ที่ช่วยให้บรรลุวัตถุประสงค์ด้านประสิทธิภาพ.....	153

สารบัญตาราง (ต่อ)

ตารางที่	หน้า
7.2 ตารางเปรียบเทียบปัจจัยการเลือกโหมดการขนส่ง.....	157
7.3 ข้อมูลพิกัด (X, Y).....	161
7.4 ตารางข้อมูลสำหรับตัวอย่างที่ 7.4.....	168
8.1 ความซับซ้อนของปัญหา VRP.....	184
8.2 คำตอบของตัวอย่างที่ 8.2.....	193
8.3 พิกัดตำแหน่งที่ตั้งของลูกค้า.....	196
8.4 ระยะทางระหว่างการเดินทางจาก โหนด i ไปยัง โหนด j.....	197
8.5 ผลประหยักระหว่างการเดินทางต่อเนื่องจาก โหนด i ไปยัง โหนด j.....	198
8.6 จัดเรียงผลประหยักระหว่างการเดินทางต่อเนื่องจาก โหนด i ไปยัง โหนด j.....	198
12.1 ผลการวิเคราะห์น้ำหนักเกณฑ์หลักและเกณฑ์รอง.....	288
12.2 คะแนนประเมินชีพพลายเออร์ทั้ง 4 ราย.....	288
13.1 เปรียบเทียบค่าพยากรณ์และค่าจริง.....	295
13.2 เปรียบเทียบค่าพยากรณ์ของแบบจำลองทั้ง 4 แบบ.....	297
13.3 เปรียบเทียบค่าพยากรณ์ของมือถือรุ่นที่ 6 ระหว่างวิธี ANN กับวิธีเดิม.....	298
13.4 การวางแผนการผลิตเริ่มตั้งแต่สัปดาห์ที่ 37 ไปจนถึงสัปดาห์ที่ 9 ของปีปัจจุบัน.....	300

สารบัญภาพ

ภาพที่	หน้า
1.1 โมเดลการบริการคลาวด์คอมพิวติ้ง.....	10
1.2 7Rs ของ SCM.....	11
1.3 ปรากฏการณ์แส้ผ้า (Bullwhip Effect).....	12
1.4 หลักการ SMART.....	15
2.1 การขยายผลกระทบของปรากฏการณ์แส้ผ้า.....	24
3.1 จำนวนลูกค้าที่ Subscribe บริการ SaaS Business ราคา \$50 ต่อเดือน.....	49
3.2 รายรับสุทธิต่อเดือน.....	50
3.3 รายรับสะสมตลอดช่วงชีวิต.....	50
3.4 รายรับตลอดช่วงชีวิต.....	51
3.5 มูลค่าของลูกค้าตลอดช่วงชีวิต.....	51
3.6 ความหมายของ RFM.....	53
4.1 กระบวนการคัดเลือกซัพพลายเออร์.....	73
5.1 องค์ประกอบของอนุกรมเวลา.....	106
6.1 ความสัมพันธ์เชิงเส้นตรง.....	126
6.2 ตัวอย่างปัญหา Heteroscedasticity.....	128
6.3 Q-Q Plot.....	129
6.4 Histogram.....	129
6.5 Q-Q Plot และ Histogram ของค่าคลาดเคลื่อน.....	132
6.6 แบบจำลอง ANN ของตัวอย่างที่ 7.3.....	141
7.1 การรวมศูนย์และการกระจายศูนย์.....	164
7.2 ผลลัพธ์จากการแก้ปัญหา LP โดยใช้โปรแกรม Python.....	169
7.3 โมเดลเชิงเส้นของตัวอย่างที่ 7.5.....	170
7.4 การป้อนข้อมูลลงใน Excel Solver ตัวอย่างที่ 7.5.....	171
8.1 ตัวอย่างเส้นทางการขนส่งอย่างง่าย.....	186
8.2 จัดเส้นทางแล้วจึงจัดกลุ่ม; จัดกลุ่มแล้วจึงจัดเส้นทาง.....	187
8.3 พิกัดตำแหน่งที่ตั้งของลูกค้า.....	196
8.4 ผลลัพธ์การจัดเส้นทาง.....	200

สารบัญภาพ (ต่อ)

ภาพที่	หน้า
8.5 พิกัดของมหาวิทยาลัยที่จะต้องไปส่งอาหาร.....	203
8.6 กลุ่มเส้นทางหลังจากทำ Spectral Clustering.....	205
8.7 เส้นทางขนส่ง.....	206
9.1 การจัดกลุ่ม RFM.....	224
9.2 การหาจำนวน cluster ที่เหมาะสมด้วยวิธี Elbow.....	227
9.3 การแจกแจงจำนวนลูกค้าในแต่ละ cluster.....	230
9.4 แผนภาพ Heat Map แสดงการทับซ้อนระหว่าง วิธี RFM กับ K-means clustering.....	232
10.1 การเปรียบเทียบระหว่างค่าจริงกับค่าพยากรณ์.....	247
10.2 การทำ Feature Importance เพื่อจัดลำดับความสำคัญของ Feature.....	250
11.1 ตำแหน่งของ Depot (★) กับ จุดส่งของ (●).....	261
11.2 การจัดกลุ่มด้วย Spectral Clustering แบ่งตามสีทั้ง 4 สี.....	262
11.3 การจัดเส้นทางขนส่งแบ่งเป็น 4 เส้นทาง ตามสีทั้ง 4 สี.....	269
12.1 แผนภูมิลำดับชั้นเชิงวิเคราะห์ในการคัดเลือกผู้ขายอย่างยั่งยืน.....	280
13.1 กราฟข้อมูลยอดขายรายสัปดาห์ของโทรศัพท์เคลื่อนที่จากผลิตภัณฑ์รุ่นก่อนหน้า.....	294
จำนวน 6 รุ่นผลิตภัณฑ์	
13.2 ขั้นตอนการพัฒนาโมเดลพยากรณ์.....	296



ส่วนที่ 1: พื้นฐานการวิเคราะห์ข้อมูลซัพพลายเชน

Part 1: Foundation of Supply Chain Data Analytics

ความสามารถในการวิเคราะห์ข้อมูล คือสมรรถนะสำคัญ
ที่ขับเคลื่อนซัพพลายเชนในยุคดิจิทัล



บทที่ 1

หลักการวิเคราะห์ข้อมูลซัพพลายเชน Supply Chain Data Analytics Principles

หัวข้อ

- 1.1 วิทยาการข้อมูล
- 1.2 ความเกี่ยวข้องกันระหว่างวิทยาการข้อมูลและคลาวด์คอมพิวติ้ง
- 1.3 การจัดการซัพพลายเชน
- 1.4 องค์ประกอบของการวิเคราะห์ข้อมูลซัพพลายเชน
- 1.5 หลักการวิเคราะห์ซัพพลายเชนอย่างมีประสิทธิภาพ

การวิเคราะห์ข้อมูล (data analytics) เป็นการจัดการข้อมูลด้วยวิธีต่าง ๆ เช่น การคำนวณ การนำเสนอข้อมูล เพื่อให้ได้ผลลัพธ์ตามวัตถุประสงค์ ซึ่งการวิเคราะห์จะเป็นการแยกแยะสิ่งที่จะพิจารณาออกเป็นส่วนย่อยที่มีความสัมพันธ์กัน เพื่อทำความเข้าใจกับข้อมูลหรือปัญหาแต่ละส่วนให้รู้ชัด รวมทั้งการสืบค้นความสัมพันธ์ของส่วนต่าง ๆ เพื่อดูว่าส่วนประกอบปลีกย่อยนั้นสามารถเข้ากันได้หรือไม่และสัมพันธ์เกี่ยวเนื่องกันอย่างไร ซึ่งจะช่วยให้เกิดความเข้าใจต่อสิ่งหนึ่งสิ่งใดอย่างแท้จริง นอกจากนี้ การวิเคราะห์ข้อมูลยังเป็นกระบวนการแปลงค่าข้อมูลให้อยู่ในรูปของผลลัพธ์ และนำผลลัพธ์ดังกล่าวมาตีความเพื่อหาข้อสรุปหรือคำตอบตามความเป็นจริง ที่สอดคล้องกับโจทย์ปัญหาหรือโจทย์วิจัยที่ตั้งไว้ การวิเคราะห์ข้อมูลมีลักษณะเป็น การลดทอน สกัด หรือกลั่นกรองข้อมูลให้อยู่ในรูปขององค์ความรู้ (knowledge) ที่สามารถอธิบายความ ทำนายปรากฏการณ์ที่ง่ายต่อการทำความเข้าใจ โดยที่การลดทอนข้อมูลจะเป็นการลดปริมาณและขนาดของข้อมูล ที่มีจำนวนมาก ๆ ให้มีขนาดเล็กลง กระทั่งเหลือข้อมูลเฉพาะที่สามารถใช้เป็นตัวแทนแห่งความคิด ในรูปของการสรุปผลในลักษณะต่าง ๆ เพียงไม่กี่ชิ้น ในปริมาณไม่มากที่จะเป็นองค์ความรู้ที่เป็นข้อสรุป (Dhar, 2013; Kudyba, 2014; Xia & Gong, 2014)

สำหรับการวิเคราะห์ข้อมูลซัพพลายเชนนั้นจะเจาะจงลงไปข้อมูลที่ไม่ใช่ข้อมูลทั่วไป แต่เป็นข้อมูลที่เกิดขึ้นในซัพพลายเชน ซึ่งเราสามารถแบ่งออกเป็นข้อมูลประเภทหลัก ๆ ได้แก่ (1) ข้อมูลที่

เกี่ยวข้องกับลูกค้า (การแบ่งกลุ่มลูกค้า หรือ ความต้องการของสินค้า) (2) ข้อมูลที่เกี่ยวข้องกับซัพพลาย (การประเมินซัพพลายเออร์ การประเมินความเสี่ยงของซัพพลาย) (3) ข้อมูลที่เกี่ยวข้องกับการจัดการคลังสินค้าและสินค้าคงคลัง (การพยากรณ์ความต้องการสินค้า หรือ การแบ่งกลุ่มสินค้าคงคลัง) และ (4) ข้อมูลเกี่ยวกับระบบโลจิสติกส์ (การออกแบบโครงข่ายโลจิสติกส์ หรือ การจัดเส้นทางการขนส่ง เป็นต้น) (Souza, 2014; Liu, 2022)

1.1 วิทยาการข้อมูล

การวิเคราะห์ข้อมูลโลจิสติกส์และซัพพลายเชนจำเป็นที่จะต้องใช้วิทยาการข้อมูล (data science) เป็นเครื่องมือประกอบในการวิเคราะห์ข้อมูลอย่างเป็นขั้นตอนและมีประสิทธิภาพโดย วิทยาการข้อมูลเป็นสหสาขาวิชาที่ใช้กระบวนการอัลกอริทึมและระบบทางวิทยาศาสตร์มาใช้ในการหาความรู้จากข้อมูลหลากหลายรูปแบบ ทั้งจัดเก็บเป็นระเบียบและไม่เป็นระเบียบ เป็นสาขาที่เกี่ยวข้องกับการทำเหมืองข้อมูล การเรียนรู้เชิงลึก และข้อมูลขนาดใหญ่ วิทยาการข้อมูลเป็นศาสตร์ที่เป็น การบูรณาการสถิติศาสตร์ การวิเคราะห์ข้อมูล และการเรียนรู้ของเครื่องเข้าด้วยกันเพื่อให้สามารถ เข้าใจและวิเคราะห์ปรากฏการณ์ที่เกิดขึ้นจริงในข้อมูลได้ โดยใช้เทคนิคและทฤษฎีที่ได้มาจาก คณิตศาสตร์ สถิติศาสตร์ วิทยาการคอมพิวเตอร์ และวิทยาการสารสนเทศ (Cao, 2017) วิทยาการ ข้อมูลมีองค์ประกอบหลักแบ่งออกเป็น 6 องค์ประกอบสำคัญ ดังนี้

1.1.1 การกำหนดปัญหา วัตถุประสงค์และคำถามงานวิจัย

การกำหนดปัญหา วัตถุประสงค์และคำถามการวิจัย ในขั้นแรกจะต้องมีการระบุปัญหาที่ต้องการตอบหรือแก้ไขอย่างชัดเจน เช่น การวิเคราะห์พฤติกรรมลูกค้า การคาดการณ์ยอดขาย หรือ การระบุปัจจัยที่ส่งผลต่อประสิทธิภาพการทำงาน เป็นต้น เมื่อทราบปัญหางานวิจัยแล้ว จะต้องมีการ ตั้งคำถาม เพื่อใช้ในการกำหนดวัตถุประสงค์ที่ชัดเจน ขั้นตอนแรกนี้ถือว่าเป็นจุดเริ่มต้นที่มีความสำคัญเพราะถ้ากำหนดวัตถุประสงค์ได้ไม่ดีพอก็จะไม่สามารถแก้ปัญหาได้ตรงจุดได้

1.1.2 การรวบรวมและจัดการข้อมูล

การรวบรวมข้อมูลจากแหล่งต่าง ๆ ซึ่งอาจเป็นข้อมูลภายในองค์กรหรือจากแหล่งภายนอก เช่น ข้อมูลจากสื่อสังคมออนไลน์ (social media) หรือข้อมูลจากแหล่งข้อมูลสาธารณะ ข้อมูลนี้อาจมาในรูปแบบที่มีโครงสร้าง เช่น ฐานข้อมูล หรือไม่มีโครงสร้าง เช่น ข้อความ รูปภาพ หรือวิดีโอ หลังจากนั้นจึงทำการรวบรวมข้อมูลจากแหล่งข้อมูลดังกล่าว ซึ่งอาจทำได้ด้วยวิธีการต่าง ๆ

เช่น การเชื่อมต่อกับฐานข้อมูลโดยตรง การใช้ส่วนต่อประสานโปรแกรมประยุกต์ (application programming interface, API) หรือการดึงข้อมูลด้วยเทคนิคการใช้โปรแกรมหรือบอท (bot) ดึงข้อมูลจากเว็บไซต์ต่าง ๆ โดยอัตโนมัติ (web scraping) แล้วจึงทำการรวมข้อมูลจากแหล่งต่าง ๆ เข้าด้วยกัน เพื่อให้สามารถนำไปใช้ในการวิเคราะห์ต่อไปได้ ซึ่งอาจต้องมีการแปลงรูปแบบข้อมูลให้สอดคล้องกันในขั้นตอนการจัดการข้อมูล

การจัดการข้อมูลรวมถึงการทำความสะอาดข้อมูล (data cleaning) คือ การจัดเก็บข้อมูล (data storage) และการเตรียมข้อมูล (data preparation) เพื่อให้พร้อมสำหรับการวิเคราะห์ข้อมูล เพราะข้อมูลที่ได้มาอาจมีข้อผิดพลาดหรือเป็นข้อมูลที่ไม่สมบูรณ์ เช่น ข้อมูลซ้ำ ข้อมูลขาดหาย หรือข้อมูลที่ผิด ดังนั้นจำเป็นต้องมีการทำความสะอาดข้อมูลเพื่อให้ได้ข้อมูลที่ถูกต้องและเหมาะสมสำหรับหลังจากนั้นจึงทำการแปลงข้อมูล ในบางครั้งข้อมูลต้องถูกแปลงให้อยู่ในรูปแบบที่เหมาะสม เช่น การแปลงประเภทของข้อมูล การปรับขนาด (scaling) การเข้ารหัส (encoding) ข้อมูลที่เป็นหมวดหมู่ (categorical data) และอื่น ๆ แล้วจึงจัดเก็บข้อมูลในระบบจัดเก็บข้อมูล เช่น ฐานข้อมูล (databases) คลังข้อมูล (data warehouses) หรือ ทะเลสาบข้อมูล (data lakes) ซึ่งต้องเลือกใช้วิธีการจัดเก็บที่เหมาะสมตามปริมาณและประเภทของข้อมูล แต่เนื่องจากในปัจจุบันข้อมูลจะมีขนาดใหญ่มาก (big data) จึงจำเป็นต้องมีการใช้เทคโนโลยีเฉพาะทาง เช่น Hadoop (ระบบจัดเก็บข้อมูลที่กระจายไฟล์ไปยังเครื่องคอมพิวเตอร์หลายเครื่องซึ่งทำงานร่วมกัน ทำให้เก็บข้อมูลได้ปริมาณมหาศาล) Spark (เครื่องมือจัดเก็บและวิเคราะห์ข้อมูลขนาดใหญ่แบบเรียลไทม์ ทำให้มีความเร็วในการประมวลผลที่รวดเร็วกว่า Hadoop) หรือ NoSQL databases (ฐานข้อมูลที่ถูกออกแบบมาเพื่อจัดเก็บข้อมูลที่ไม่มีโครงสร้างชัดเจน เช่น รูปภาพ วิดีโอ ข้อความเสียง และมีความยืดหยุ่นสูง ขยายขนาดระบบจัดเก็บได้ง่ายจึงเหมาะกับข้อมูลที่หลากหลายและมีปริมาณมหาศาลได้อย่างรวดเร็ว)

โดยสรุปการรวบรวมและจัดการข้อมูลเป็นขั้นตอนสำคัญที่มีผลต่อความสำเร็จของการวิเคราะห์ข้อมูลและการสร้างโมเดลในวิทยาการข้อมูล การทำงานอย่างละเอียดในขั้นตอนนี้ช่วยให้มั่นใจได้ว่าข้อมูลที่ใช้ในการวิเคราะห์นั้นมีคุณภาพสูงและเหมาะสมกับการใช้งาน

1.1.3 การวิเคราะห์ข้อมูล

การใช้เทคนิคทางสถิติและวิธีการเชิงปริมาณในการสำรวจและวิเคราะห์ข้อมูล เพื่อค้นหารูปแบบ แนวโน้ม หรือความสัมพันธ์ที่ซ่อนอยู่ในข้อมูล การวิเคราะห์ข้อมูลอาจใช้เครื่องมือและซอฟต์แวร์ต่าง ๆ เช่น Python R SQL Excel และเครื่องมือทางสถิติอื่น ๆ หรืออีกนัยหนึ่ง การวิเคราะห์ข้อมูล คือกระบวนการตรวจสอบ แปลง และสร้างแบบจำลองข้อมูลเพื่อค้นหาข้อมูลเชิงลึก ระบุรูปแบบ หรือทำการสรุปผลที่มีความหมาย ซึ่งสามารถนำไปใช้ในการตัดสินใจหรือแก้ไขปัญหาในธุรกิจและการวิจัยได้ (Provost & Fawcett, 2013) การวิเคราะห์ข้อมูลสามารถจำแนกออกเป็นประเภทหลัก ๆ ได้แก่

1. การวิเคราะห์เชิงพรรณนา (Descriptive Analysis) เป็นการสรุปข้อมูลเพื่อให้เข้าใจอดีตหรือสถานการณ์ปัจจุบันเครื่องมือที่ใช้บ่อย เช่น การคำนวณค่าเฉลี่ย (mean) ตารางความถี่ (frequency tables) ส่วนเบี่ยงเบนมาตรฐาน (standard deviation) และการสร้างกราฟต่าง ๆ (Han et al., 2011)

2. การวิเคราะห์เชิงวินิจฉัย (Diagnostic Analysis) มุ่งเน้นไปที่การทำความเข้าใจสาเหตุหรือปัจจัยที่นำไปสู่เหตุการณ์บางอย่างวิธีการนี้อาจใช้การวิเคราะห์ความสัมพันธ์ (correlation analysis) หรือการถดถอย (regression analysis) เพื่อหาเหตุผลและสาเหตุของปัญหา (Han et al., 2011)

3. การวิเคราะห์เชิงทำนาย (Predictive Analysis) การใช้ข้อมูลในอดีตและปัจจุบันเพื่อทำนายผลลัพธ์ในอนาคต เทคนิคนี้รวมถึงการสร้างโมเดลทำนายด้วยการเรียนรู้ของเครื่อง (machine learning) การวิเคราะห์อนุกรมเวลา (time series analysis) และการทำเหมืองข้อมูล (data mining) เป็นต้น (Han et al., 2011)

4. การวิเคราะห์เชิงสั่งการ (Prescriptive Analysis) มุ่งเน้นที่การกำหนดวิธีการหรือแนะนำการกระทำเพื่อให้ได้ผลลัพธ์ที่ดีที่สุด ใช้การจำลองสถานการณ์ (simulation) การเพิ่มประสิทธิภาพ (optimization) และการวิเคราะห์เชิงคาดการณ์แบบซับซ้อนเพื่อแนะนำแนวทางปฏิบัติ (Han et al., 2011)

สำหรับขั้นตอนการวิเคราะห์ข้อมูลจะประกอบด้วย 4 ขั้นตอนได้แก่ (1) การสำรวจและวิเคราะห์ข้อมูลเบื้องต้น คือการสำรวจข้อมูลเบื้องต้นเพื่อให้เข้าใจโครงสร้างของข้อมูลและตรวจสอบความสัมพันธ์ระหว่างตัวแปรต่าง ๆ (2) การสร้างแบบจำลอง เป็นการใช้โมเดลสถิติหรือการใช้เทคนิค

การเรียนรู้ของเครื่อง เพื่อวิเคราะห์ข้อมูลและทำนายผลลัพธ์ (3) การแปลผลและสรุปผล คือการนำผลลัพธ์ที่ได้มาวิเคราะห์เพื่อสรุปเป็นข้อมูลเชิงลึกและนำเสนอให้กับผู้มีส่วนเกี่ยวข้อง (4) การนำไปใช้และติดตามผล คือการนำผลการวิเคราะห์ไปใช้ในการตัดสินใจหรือแก้ไขปัญหาและทำการติดตามผลลัพธ์เพื่อปรับปรุงวิธีการหากจำเป็น

โดยสรุปการวิเคราะห์ข้อมูลเป็นกระบวนการที่จำเป็นสำหรับองค์กรและธุรกิจเพื่อให้สามารถตัดสินใจได้อย่างมีข้อมูลรองรับ และยังช่วยในการเพิ่มประสิทธิภาพการดำเนินการให้กับองค์กรเพื่อให้ประสบความสำเร็จตามเป้าหมายที่ตั้งไว้

1.1.4 การเรียนรู้ของเครื่อง (Machine Learning: ML)

การเรียนรู้ของเครื่องเป็นหัวใจสำคัญของวิทยาการข้อมูลและเป็นสาขาย่อยของปัญญาประดิษฐ์ (artificial intelligence หรือ AI) ที่เน้นการพัฒนาอัลกอริทึมและโมเดลที่สามารถเรียนรู้จากข้อมูลและทำการคาดการณ์หรือการตัดสินใจได้โดยไม่ต้องมีการเขียนโปรแกรมเฉพาะเจาะจงสำหรับทุกสถานการณ์ (Mitchell, 1997; Hastie et al., 2009)

หลักการสำคัญของการเรียนรู้ของเครื่องแบ่งออกเป็น 3 ข้อได้แก่ (1) การเรียนรู้ของเครื่องเกี่ยวข้องกับการให้คอมพิวเตอร์วิเคราะห์ข้อมูลจำนวนมากเพื่อค้นหารูปแบบ (patterns) และความสัมพันธ์ระหว่างตัวแปรต่าง ๆ จากนั้นจึงสร้างแบบจำลองที่สามารถทำการคาดการณ์หรือตัดสินใจได้โดยอัตโนมัติ (2) การทำซ้ำและการปรับปรุง คือการใช้อัลกอริทึมในการเรียนรู้ของเครื่องมาทำการปรับปรุงตัวเองโดยการเรียนรู้จากข้อผิดพลาดและปรับโมเดลเพื่อให้มีความแม่นยำสูงขึ้นในแต่ละรอบของการฝึกอบรม (3) การใช้งานในสถานการณ์ที่หลากหลาย เพราะการเรียนรู้ของเครื่องถูกนำมาใช้ในหลากหลายสถานการณ์ เช่น การคาดการณ์ยอดขาย การจดจำภาพ การวิเคราะห์ข้อความ การจำแนกประเภท (classification) และการจัดกลุ่ม (clustering) (Mitchell, 1997; Hastie et al., 2009)

ประเภทของการเรียนรู้ของเครื่องสามารถแบ่งออกได้เป็น 5 ประเภทได้แก่ (1) การเรียนรู้ภายใต้การควบคุม (supervised learning) คือการใช้ข้อมูลที่มีป้ายกำกับ (labeled data) ซึ่งประกอบด้วยคู่ของอินพุตและผลลัพธ์ที่ต้องการ โมเดลจะเรียนรู้จากข้อมูลเหล่านี้และสามารถทำนายผลลัพธ์ของข้อมูลใหม่ได้ ตัวอย่างของอัลกอริทึมได้แก่ การถดถอยเชิงเส้น (linear regression) การจำแนกประเภทด้วยต้นไม้ตัดสินใจ (decision trees) และการจำแนกประเภทด้วยเครื่องเวกเตอร์

สนับสนุน (support vector machines) (2) การเรียนรู้โดยไม่มีการควบคุม (unsupervised learning) คือการใช้ข้อมูลที่ไม่มีป้ายกำกับ โดยอัลกอริทึมจะค้นหารูปแบบและโครงสร้างในข้อมูล โดยอัตโนมัติ ตัวอย่างของอัลกอริทึมได้แก่ การจัดกลุ่ม (clustering) เช่น K-Means การลดมิติ (dimensionality reduction) เช่น PCA (3) การเรียนรู้แบบเสริมกำลัง (reinforcement learning) เกี่ยวข้องกับการฝึกอบรมตัวแทน (agent) ให้ตัดสินใจในสภาพแวดล้อมโดยการเรียนรู้จากผลตอบรับ (rewards) และการลงโทษ (punishments) ตัวอย่างของอัลกอริทึมได้แก่ Q-Learning และ Deep Q-Networks (DQN) (Sutton & Barto, 2018) (4) การเรียนรู้กึ่งควบคุม (semi-supervised learning) คือ การผสมผสานการเรียนรู้ภายใต้การควบคุมและการเรียนรู้โดยไม่มีการควบคุม โดยใช้ข้อมูลที่มีป้ายกำกับจำนวนน้อยและข้อมูลที่ไม่มีป้ายกำกับจำนวนมากเพื่อปรับปรุงประสิทธิภาพของโมเดล (Zhu, 2005) (5) การเรียนรู้แบบโอนย้าย (transfer learning) คือ การใช้ความรู้ที่ได้จากงานหนึ่งเพื่อปรับปรุงประสิทธิภาพของโมเดลในงานอื่นที่มีความคล้ายคลึงกัน (Pan & Yang, 2010)

ในปัจจุบันการเรียนรู้ของเครื่อง ได้กลายเป็นเครื่องมือที่ทรงพลังในการขับเคลื่อนนวัตกรรมและความก้าวหน้าในหลาย ๆ สาขา รวมถึงการเงิน การแพทย์ การตลาด และวิทยาการคอมพิวเตอร์

1.1.5 การแสดงผลข้อมูล

การนำข้อมูลและผลการวิเคราะห์มานำเสนอในรูปแบบกราฟ แผนภูมิ แผนภาพ หรือการแสดงผลเชิงภาพอื่น ๆ จะช่วยให้ผู้ใช้สามารถเข้าใจข้อมูลเชิงลึกได้อย่างง่ายดายและเห็นภาพรวมของข้อมูลได้ชัดเจน (Munzner, 2014) โดยในเบื้องต้นเราสามารถแสดงผลข้อมูลในรูปแบบกราฟประเภทต่าง ๆ ได้แก่กราฟแท่ง (bar charts) ใช้เพื่อเปรียบเทียบค่าของตัวแปรต่าง ๆ โดยใช้แท่งที่มีความยาวแตกต่างกันเพื่อแสดงขนาดของข้อมูล กราฟเส้น (line charts) ใช้เพื่อแสดงแนวโน้มของข้อมูลตลอดเวลาหรือการเปลี่ยนแปลงของข้อมูลในลำดับที่เป็นต่อเนื่อง กราฟวงกลม (pie charts) ใช้เพื่อแสดงการแบ่งสัดส่วนของข้อมูลที่เป็นส่วนหนึ่งของทั้งหมด โดยแสดงข้อมูลเป็นชิ้นส่วนของวงกลม แผนที่ความร้อน (heat maps) ใช้เพื่อแสดงข้อมูลในรูปแบบของตาราง โดยใช้สีเพื่อแสดงความหนาแน่นหรือระดับของข้อมูล กราฟจุด (scatter plots) ใช้เพื่อแสดงความสัมพันธ์ระหว่างตัวแปรสองตัวโดยใช้จุดที่กระจายอยู่ในกราฟ กราฟกล่อง (box plots) ใช้เพื่อแสดงการกระจายของข้อมูลและตรวจสอบข้อมูลที่เป็นข้อยกเว้น (outliers) กราฟพื้นที่ (area charts) คล้าย

กับกราฟเส้น แต่จะเติมสีพื้นที่ด้านล่างของเส้นเพื่อแสดงการเปลี่ยนแปลงของข้อมูล กราฟสามมิติ (3D charts) ใช้เพื่อแสดงข้อมูลในมิติที่สาม เพิ่มความลึกให้กับการแสดงผล เช่น กราฟแท่งสามมิติ (3D bar charts) แผนที่ (maps) ใช้เพื่อแสดงข้อมูลที่มีความสัมพันธ์กับตำแหน่งทางภูมิศาสตร์ เช่น แผนที่การกระจายของประชากร หรือแผนที่ความร้อน โดยสรุปแล้วการแสดงผลข้อมูลเป็นเครื่องมือที่สำคัญในการทำให้ข้อมูลซับซ้อนกลายเป็นสิ่งที่สามารถเข้าใจได้ง่ายและนำไปใช้ในการตัดสินใจได้อย่างมีประสิทธิภาพ

1.1.6 การสื่อสารและการตัดสินใจ

การสื่อสารผลการวิเคราะห์ข้อมูลและข้อมูลเชิงลึกให้กับผู้มีส่วนเกี่ยวข้อง เช่น ผู้บริหาร หรือผู้ใช้ในธุรกิจ เพื่อช่วยในการตัดสินใจที่มีข้อมูลรองรับ โดยที่นักวิทยาการข้อมูลต้องสามารถอธิบายผลลัพธ์ที่ซับซ้อนในรูปแบบที่ง่ายต่อการเข้าใจสำหรับผู้ฟังที่ไม่มีพื้นฐานด้านเทคนิค (Knaflic, 2015) การสื่อสารถึงแม้เป็นส่วนประกอบสุดท้ายของวิทยาการข้อมูลแต่ก็ถือว่าเป็นส่วนที่สำคัญไม่ยิ่งหย่อนไปกว่าส่วนประกอบอื่น ๆ เพราะถึงแม้จะมีการรวบรวมข้อมูลที่ดี การวิเคราะห์ที่สามารถทำนายผลลัพธ์ในอนาคตได้แม่นยำเพียงใด แต่ถ้าไม่สามารถที่จะอธิบายให้ผู้ฟัง ผู้ใช้งาน หรือผู้บริหารให้เข้าใจเพื่อใช้ประกอบการตัดสินใจได้อย่างมีประสิทธิภาพ ทุกอย่างที่ทำมาก็ไร้ประโยชน์ ดังนั้นจึงจำเป็นอย่างยิ่งที่จะต้องใช้เทคนิคต่าง ๆ ประกอบ ยกตัวอย่าง เช่น (1) การสื่อสารที่ชัดเจน และกระชับ ใช้ภาษาที่เข้าใจง่าย หลีกเลี่ยงการใช้คำศัพท์ที่ซับซ้อนหรือเทคนิคเกินไป ตรงประเด็นไม่กล่าวถึงประเด็นหลักและหลีกเลี่ยงการเบี่ยงเบนจากหัวข้อหลัก (2) การใช้ภาษากายและน้ำเสียง ใช้ภาษากายที่เป็นมิตรและเปิดเผย เช่น การสบตาและท่าทางเป็นมิตร น้ำเสียง ใช้น้ำเสียงที่เหมาะสม และสื่อถึงความตั้งใจหรืออารมณ์ที่ต้องการ (3) การตั้งคำถามที่เหมาะสม โดยใช้คำถามเปิด ใช้คำถามเปิดที่ช่วยกระตุ้นการสนทนา เช่น “คุณคิดอย่างไรเกี่ยวกับ...?” หรือคำถามปิดโดยใช้คำถามปิดเพื่อการตอบที่เฉพาะเจาะจง เช่น “คุณมีเวลาในวันพรุ่งนี้ไหม?” (4) การใช้สื่อที่เหมาะสม โดยใช้ช่องทางการสื่อสารที่เหมาะสมกับเนื้อหาหรือผู้รับสาร เช่น การประชุมตัวต่อตัวสำหรับการสนทนาที่ซับซ้อน หรืออีเมลสำหรับการสื่อสารที่เป็นทางการเป็นต้น (5) การสร้างความสัมพันธ์ที่ดีโดยใช้การสร้างความสัมพันธ์ที่ดีเพื่อให้การสื่อสารเป็นไปอย่างราบรื่นแสดงความเข้าใจและเห็นอกเห็นใจต่อความรู้สึกของผู้อื่น (6) การใช้ตัวอย่างและการเปรียบเทียบโดยใช้ตัวอย่างจริงหรือเปรียบเทียบเพื่ออธิบายแนวคิดที่ซับซ้อนให้เข้าใจง่ายสร้างภาพจินตนาการ

โดยสรุปการรวมกันของส่วนประกอบเหล่านี้ทำให้วิทยาการข้อมูลเป็นสาขาที่มีความหลากหลายและมีบทบาทสำคัญในการนำข้อมูลไปใช้ประโยชน์ในหลาย ๆ ด้าน

1.2 ความเกี่ยวข้องกันระหว่างวิทยาการข้อมูลและคลาวด์คอมพิวติ้ง

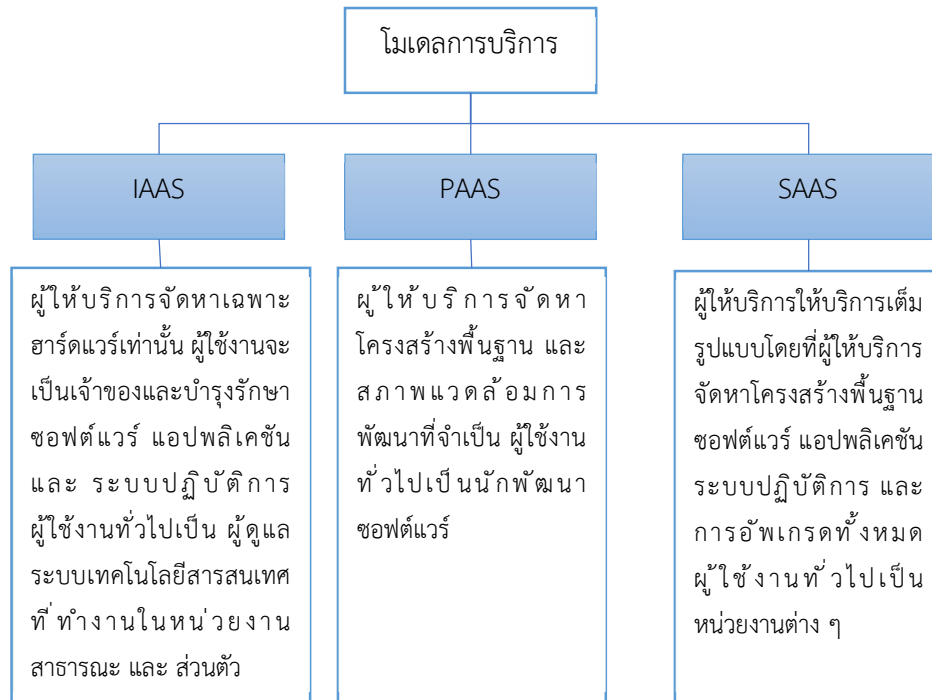
วิทยาการข้อมูล และ คลาวด์คอมพิวติ้ง (cloud computing) เป็นสองสาขาที่มีความเกี่ยวข้องกันอย่างมากในปัจจุบัน โดยทั้งสองสามารถทำงานร่วมกันเพื่อสร้างประโยชน์ในหลาย ๆ ด้าน โดยเฉพาะในบริบทของการจัดการและวิเคราะห์ข้อมูลขนาดใหญ่รวมถึงการทำงานที่ต้องการทรัพยากรการประมวลผลสูง ในด้านการจัดเก็บข้อมูลขนาดใหญ่ คลาวด์คอมพิวติ้ง ช่วยให้สามารถจัดเก็บข้อมูลขนาดใหญ่ได้อย่างมีประสิทธิภาพ โดยการใช้บริการที่มีความยืดหยุ่นในการจัดเก็บและเข้าถึงข้อมูล ทำให้นักวิทยาศาสตร์ข้อมูลสามารถเข้าถึงข้อมูลที่ต้องการได้จากทุกที่ทุกเวลา การจัดเก็บข้อมูลบนคลาวด์ช่วยลดต้นทุนและเพิ่มความยืดหยุ่นเมื่อเทียบกับการใช้เซิร์ฟเวอร์ภายในองค์กร (Areerakulkan & Pongpech, 2021)

สำหรับการประมวลผลข้อมูล วิทยาการข้อมูลมักต้องการประมวลผลข้อมูลจำนวนมาก ซึ่งอาจใช้เวลานานหากใช้คอมพิวเตอร์หรือเซิร์ฟเวอร์ภายในองค์กร แต่คลาวด์คอมพิวติ้ง สามารถให้บริการประมวลผลที่รวดเร็วและมีประสิทธิภาพ ที่สามารถปรับขนาดทรัพยากรตามความต้องการได้ นอกจากนั้นการใช้คลาวด์คอมพิวติ้ง ช่วยให้สามารถประมวลผลข้อมูลขนาดใหญ่หรือทำงานด้านการเรียนรู้ของเครื่องได้รวดเร็วขึ้น (Kochan et al., 2018) การใช้แพลตฟอร์มและเครื่องมือบนคลาวด์แพลตฟอร์มและเครื่องมือบนคลาวด์ที่ถูกพัฒนาขึ้นมาเพื่อสนับสนุนงานวิทยาการข้อมูลช่วยให้นักวิทยาศาสตร์ข้อมูลสามารถทำการวิเคราะห์และสืบค้นข้อมูลได้อย่างรวดเร็วและมีประสิทธิภาพ นอกจากนี้บริการคลาวด์ยังมีเครื่องมือสำหรับการเรียนรู้ของเครื่องที่พัฒนาขึ้นช่วยให้การปรับใช้โมเดลการเรียนรู้ของเครื่องเป็นไปได้ง่ายและรวดเร็ว การใช้คลาวด์คอมพิวติ้งยังช่วยให้นักวิทยาศาสตร์ข้อมูลจากทั่วโลกสามารถทำงานร่วมกันได้อย่างมีประสิทธิภาพผ่านทางแพลตฟอร์มบนคลาวด์ที่อนุญาตให้แชร์ข้อมูล โค้ด (code) และผลลัพธ์ได้อย่างสะดวก นอกจากนี้คลาวด์คอมพิวติ้งยังช่วยให้การทำงานสามารถทำได้จากทุกที่ที่มีการเชื่อมต่ออินเทอร์เน็ต ลดข้อจำกัดเรื่องการเข้าถึงทรัพยากรการประมวลผล (Neaga et al., 2015)

ในมุมมองของการจัดการซัพพลายเชนคลาวด์คอมพิวติ้งส่งเสริมโอกาสที่สำคัญของการแชร์ข้อมูลคู่ค้าอย่างเรียลไทม์รวมถึงการตอบโต้และการร่วมมือกัน ซึ่งส่งผลอย่างสำคัญในกระบวนการ

ทำงานของซัพพลายเชน เช่น การพยากรณ์อุปสงค์ (demand) ของลูกค้าสุดท้ายสามารถที่จะทำได้เพียงครั้งเดียวสำหรับตลอดทั้งซัพพลายเชนบนพื้นฐานของข้อมูลระบบขายหน้าร้าน (point of sale, POS) หลังจากนั้นผลลัพธ์ของการพยากรณ์จะถูกแชร์ให้กับคู่ค้าทั้งหมดในซัพพลายเชน ซึ่งจะช่วยลดการขยายสัญญาณของอุปสงค์จากปรากฏการณ์แส้ม้า (bullwhip effect) (ดูภาพที่ 1.3) ทำให้การวางแผน การจัดตาราง การจัดการขนส่งและคลังสินค้า การจัดการสินค้าคงคลังตลอดซัพพลายเชนเป็นไปตามแผนอย่างมีประสิทธิภาพ

1.2.1 โครงสร้างของคลาวด์คอมพิวเตอร์ที่ใช้งานในปัจจุบัน สามารถแบ่งเป็น 3 ลักษณะ ได้แก่ (1) โครงสร้างคลาวด์สาธารณะ (public cloud structure) ผู้ใช้งานทุกคนจะแชร์คลาวด์ โดยที่ความปลอดภัยของข้อมูลผู้ใช้และการกำหนดสิทธิเข้าถึงของผู้ใช้สามารถทำได้ แต่จะไม่เข้มงวดเท่ากับคลาวด์ส่วนตัว (private cloud) (2) โครงสร้างคลาวด์แบบผสม (hybrid cloud structure) องค์กรจะเก็บรักษาข้อมูลที่สำคัญไว้ในฐานข้อมูลและสามารถเข้าถึงได้ผ่านเซิร์ฟเวอร์ภายในองค์กร ส่วนข้อมูลและแอปพลิเคชันที่ไม่เป็นความลับอาจถูกจัดเก็บ เข้าถึง และ สอบถามบนคลาวด์ได้ (3) โครงสร้างคลาวด์ส่วนตัว (private cloud structure) ผู้ให้บริการคลาวด์จัดหาทุก ๆ สิ่งที่เป็นสำหรับองค์กรแต่ละองค์กรโดยเฉพาะ ยกตัวอย่าง เช่น โรงพยาบาลเป็นเจ้าแรก ๆ ที่ใช้คลาวด์ส่วนตัวทำการเก็บข้อมูลสุขภาพคนไข้และเวชระเบียนรวมถึงข้อมูลการทำธุรกรรมของคนไข้ (วันเข้า วันออก ค่าธรรมเนียม ประกัน ใบแจ้งหนี้ ใบเสร็จ) สำหรับโมเดลการให้บริการจะถูกแบ่งออกเป็น 3 แบบ ได้แก่ (1) infrastructure as a service (IAAS) (2) platform as a service (PASS) และ (3) software as a service (SAAS) (ดูภาพที่ 1.1)



ภาพที่ 1.1 โมเดลการบริการคลาวด์คอมพิวเตอร์ ระดับเริ่มต้นที่ผู้ให้บริการจัดหาเฉพาะฮาร์ดแวร์เท่านั้น (IAAS) ระดับกลางที่ผู้ให้บริการจัดหาฮาร์ดแวร์ และ สภาพแวดล้อมที่จำเป็น (PAAS) และระดับสูงที่ผู้ให้บริการให้บริการเต็มรูปแบบ (SAAS)

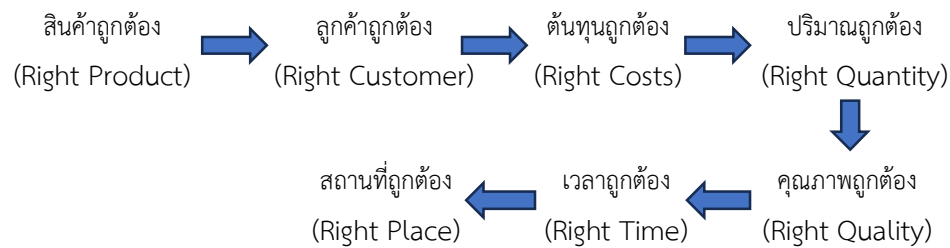
ที่มา: ภาพโดยผู้เขียน

1.3 การจัดการซัพพลายเชน (Supply Chain Management, SCM)

การจัดการซัพพลายเชนคือการบูรณาการ และการร่วมมือเชิงกลยุทธ์อย่างมีระบบของการทำงานทางธุรกิจ และ กระบวนการข้ามซัพพลายเชนเพื่อสร้างคุณค่าสูงสุดให้กับลูกค้าอย่างมีประสิทธิภาพ ซึ่งก็คือการจัดการกระบวนการทั้งภายใน และ ภายนอกซัพพลายเชนที่เกี่ยวกับการไหลความสัมพันธ์ระหว่าง ผู้ซื้อ-ผู้จัดหา (buyer-supplier) และ โครงสร้างของโครงข่ายการเชื่อมโยงทางธุรกิจรวมถึงการวางแผนกลยุทธ์ และการจัดเรียงกลยุทธ์ซัพพลายเชนให้ตรงกับเป้าหมายภาพรวมทางธุรกิจ

การจัดการซัพพลายเชนมีวัตถุประสงค์ที่จะให้บริการ และ ประสิทธิภาพกับลูกค้าอย่างยกระดับผ่านการจัดการอย่างสอดคล้องกันของการไหลของสินค้าร่วมกับข้อมูล และการเงินจากจุดกำเนิดไปยังจุดสุดท้ายของการบริโภค หรือจะพูดง่าย ๆ ก็ คือ การจัดการซัพพลายเชนมุ่งที่จะสร้าง

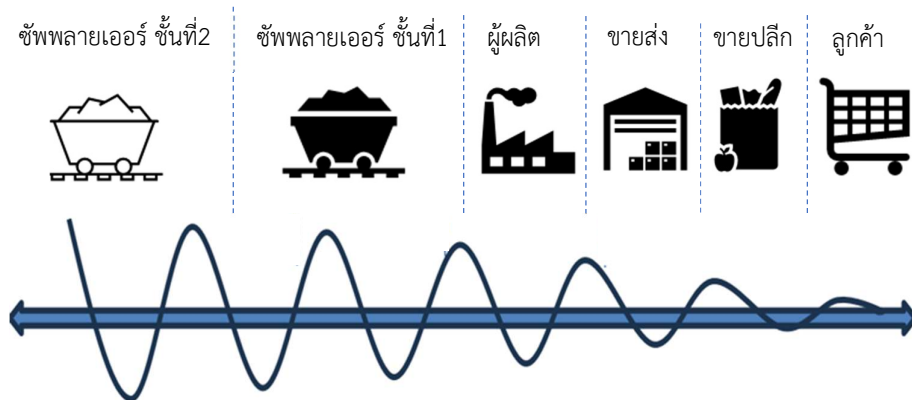
หรือผลิตสินค้าและบริการที่ถูกต้อง (right product) ให้กับลูกค้าที่ใช่ (right customer) ในต้นทุนที่เหมาะสม (right cost) ในปริมาณตามที่ต้องการ (right quantity) คุณภาพที่ใช่ (right quality) ตามเวลาที่กำหนด (right time) และ สถานที่ที่ถูกต้อง (right place) (หรือเรียกสั้น ๆ ว่า 7Rs) การจัดการซัพพลายเชนอย่างมีประสิทธิภาพควรที่จะสอดคล้องอย่างใกล้ชิดกับกลยุทธ์ทางธุรกิจโดยรวม และ เพิ่มความได้เปรียบในการแข่งขันทางธุรกิจ และ ความสามารถในการทำกำไรโดยยกระดับความพึงพอใจของลูกค้าโดยรวม (ภาพที่ 1.2)



ภาพที่ 1.2 7Rs ของ SCM

ที่มา: ดัดแปลงมาจาก (Simchi-Levi, Kaminsky, & Simchi-Levi, 2004)

จากภาพที่ 1.2 การที่จะประสบความสำเร็จใน 7Rs จำเป็นอย่างยิ่งที่องค์กรจะต้องมีการจัดการความร่วมมือในกระบวนการต่างของซัพพลายเชน อย่างไรก็ตามการจัดการความร่วมมืออย่างมีประสิทธิภาพจำเป็นจะต้องมีการแลกเปลี่ยนข้อมูลที่แม่นยำ ต่อเนื่อง และ ทันเวลา ตลอดทั้งซัพพลายเชน ถ้าปราศจากข้อมูลที่สำคัญ เช่น อุปสงค์ของลูกค้า ยอดขาย ระดับสต็อกปัจจุบัน และ กำลังการผลิตที่มีอยู่ คงเป็นเรื่องที่ยากมากสำหรับผู้ผลิตที่จะทำการตัดสินใจอย่างมีประสิทธิภาพสำหรับการวางแผนและการจัดตารางการผลิต นอกจากนั้นการแลกเปลี่ยนข้อมูลที่ไม่เที่ยงตรงอาจเป็นสาเหตุที่สำคัญของการเกิดปรากฏการณ์สั้มน้ำไปยังต้นน้ำของซัพพลายเชนซึ่งส่งผลต่อประสิทธิภาพโดยรวมและอาจก่อให้เกิดความสูญเปล่ารวมถึงต้นทุนสินค้าคงคลังที่สูงขึ้นดังแสดงในภาพที่ 1.3



ภาพที่ 1.3 ปรากฏการณ์เส้มาเป็นความผิดปกติในซัพพลายเชนที่อุปสงค์ของลูกค้าปลายทางเปลี่ยนแปลงเพียงเล็กน้อย แต่กลับถูกขยายผลให้ใหญ่ขึ้นเรื่อย ๆ เมื่อส่งผ่านจากค้าปลีกไปยังผู้ค้าส่ง ผู้ผลิต และ ซัพพลายเออร์ เปรียบเสมือนการสะบัดเส้ ที่มีอวบเพียงนิดเดียวแต่ปลายเส้กลับสะบัดอย่างรุนแรง ส่งผลให้มีปัญหาสินค้าขาดแคลนหรือ สินค้าล้นสต็อก (overstock)

ที่มา: ดัดแปลงมาจาก (Udenio et al., 2017)

ด้วยการตระหนักถึงความสำคัญของการแลกเปลี่ยนข้อมูลที่จำเป็น ในปัจจุบันจึงได้มีการใช้ซอฟต์แวร์ต่าง ๆ เพื่อช่วยในการแลกเปลี่ยนข้อมูลอย่างไร้รอยต่อ ได้แก่ระบบ ERP หรือ CRM อย่างไรก็ตามยังมีอุปสรรคที่สำคัญของการใช้งานซอฟต์แวร์เหล่านี้ ได้แก่ การที่พนักงานไม่ได้ถูกอบรมอย่างเพียงพอที่จะใช้งานระบบได้ ทำให้ไม่สามารถที่จะใช้ประโยชน์ได้อย่างเต็มที่ การมีข้อผิดพลาด (bugs) เพราะไม่ได้ปรับให้สอดคล้องกับองค์กรอย่างแท้จริงทำให้ขัดต่อการใช้งาน การขาดความเชื่อถือกันระหว่างสมาชิกของซัพพลายเชนทำให้เกิดความไม่เต็มใจที่จะแบ่งปันข้อมูลที่สำคัญ รวมไปถึงระบบที่เข้ากันไม่ได้ในซัพพลายเชน (ระบบปฏิบัติการคนละแบบที่ถูกใช้โดยซัพพลายเออร์) รวมไปถึงการขาดความร่วมมือทั้งภายในและระหว่างองค์กรทำให้เกิดอุปสรรคในการแบ่งปันข้อมูล ดังนั้นเพื่อเอาชนะอุปสรรคดังที่กล่าวมาเพื่อให้เกิดการแบ่งปันข้อมูลอย่างมีประสิทธิภาพ จำเป็นที่องค์กรจะต้องสร้างพันธมิตรทางการค้ากับสมาชิกในซัพพลายเชนเพื่อทำให้เกิดความเชื่อใจระหว่างกันส่งผลให้สมาชิกเต็มใจที่จะแบ่งปันข้อมูลที่สำคัญให้กับสมาชิกอื่น ๆ อย่างทันเวลา

ในการจัดการซัพพลายเชนข้อมูลเหล่านี้สามารถแบ่งออกเป็น 3 ชนิดพื้นฐานได้แก่ (1) ข้อมูลที่มีโครงสร้างชัดเจน (structured data) คือ ข้อมูลที่มีรูปแบบมาตรฐานที่พร้อมจะถูกเก็บ ดึง และ

ประมวลผล ตัวอย่างของข้อมูลชนิดนี้ได้แก่ Excel หรือ ฐานข้อมูล SQL ข้อมูล GPS ข้อมูลบันทึกยอดขาย ข้อมูลทางการเงินจากตลาดหุ้น (2) ข้อมูลที่ไม่มีโครงสร้าง (unstructured data) คือ ข้อมูลที่ไม่มีรูปแบบมาตรฐาน ซึ่งจะใช้ระยะเวลามากกว่าในการประมวลผล ตัวอย่างของข้อมูลชนิดนี้ได้แก่ บันทึกเสียง วีดีโอ ข้อเสนอแนะของลูกค้า หรือ โพสต์ทวิตเตอร์ เป็นต้น (3) ข้อมูลกึ่งโครงสร้าง (semi-structured data) ไม่มีนิยามที่ชัดเจนสำหรับข้อมูลชนิดนี้แต่โดยทั่วไปจะประกอบไปด้วยข้อมูลทั้งสองรูปแบบข้างต้น (มีโครงสร้างและไม่มีโครงสร้าง) ได้แก่ ไฟล์จาวก .JSON .XML และ .CSV เป็นต้น ข้อมูลกึ่งโครงสร้างจะสามารถที่จะถูกประมวลผลและวิเคราะห์ได้ง่ายกว่าข้อมูลที่ไม่มีโครงสร้าง เนื่องจากเครื่องมือในการวิเคราะห์ในปัจจุบันสามารถที่จะแปลง อ่าน และ ประมวลผลข้อมูลชนิดนี้ได้อย่างรวดเร็ว

1.4 องค์ประกอบของการวิเคราะห์ข้อมูลซัพพลายเชน

การวิเคราะห์ข้อมูลซัพพลายเชนเป็นกระบวนการที่ใช้ข้อมูลเพื่อทำความเข้าใจและปรับปรุงประสิทธิภาพของซัพพลายเชน (Waller & Fawcett, 2013; Souza, 2014) ซึ่งประกอบด้วยหลายส่วนสำคัญที่ช่วยให้การวิเคราะห์มีประสิทธิภาพและสามารถตัดสินใจโดยใช้ข้อมูลพื้นฐานประกอบการตัดสินใจ องค์ประกอบในการวิเคราะห์ข้อมูลซัพพลายเชนจะมีลักษณะเดียวกันกับขั้นตอนการวิเคราะห์ข้อมูลที่ได้อธิบายไปเบื้องต้น ซึ่งสามารถอธิบายโดยสรุปได้ดังนี้

1. การรวบรวมข้อมูลเกิดขึ้นในซัพพลายเชนที่เกี่ยวข้องหรือต้องการจะนำมาวิเคราะห์ เช่น ข้อมูลทางการเงิน (ข้อมูลเกี่ยวกับต้นทุนการผลิต ต้นทุนการจัดเก็บ ต้นทุนการขนส่ง และรายได้จากการขาย) ข้อมูลการสั่งซื้อ (ข้อมูลเกี่ยวกับคำสั่งซื้อ ปริมาณ และเวลาที่ส่งมอบ) ข้อมูลการผลิต (ข้อมูลเกี่ยวกับกระบวนการผลิต คุณภาพสินค้า และเวลาการผลิต) ข้อมูลการขนส่ง (ข้อมูลเกี่ยวกับวิธีการขนส่ง เวลาการจัดส่ง และความสามารถในการขนส่ง) หรือ ข้อมูลการจัดเก็บ (ข้อมูลเกี่ยวกับระดับสต็อก พื้นที่จัดเก็บ และการเคลื่อนไหวของสินค้าคงคลัง)

2. การจัดการข้อมูลเริ่มจากการทำความสะอาดข้อมูล โดยทำการตรวจสอบและแก้ไขข้อผิดพลาดในข้อมูล เช่น ข้อมูลที่หายไปหรือข้อมูลที่ไม่ครบหรือไม่ถูกต้อง หลังจากนั้นจึงทำการจัดเก็บข้อมูลที่ผ่านมาทำความสะอาดแล้วในฐานข้อมูลหรือระบบคลาวด์ที่สามารถเข้าถึงและจัดการได้อย่างมีประสิทธิภาพ เพื่อจะนำมาเข้าสู่การวิเคราะห์ข้อมูลในขั้นตอนถัดไป

3. การวิเคราะห์ข้อมูลสามารถทำได้ในหลายรูปแบบตามลำดับ โดยเริ่มจากการวิเคราะห์ข้อมูลเพื่อทำความเข้าใจภาพรวมของข้อมูลและค้นหาสาเหตุของปัญหาที่เกิดขึ้นในปัจจุบัน หรือที่เรียกว่าการวิเคราะห์เชิงพรรณนาและการวิเคราะห์หาสาเหตุ (descriptive and diagnostic analytics) โดยสามารถใช้เครื่องมือทางสถิติหรือการเรียนรู้ของเครื่องเข้ามาประกอบการวิเคราะห์ ตัวอย่างของการวิเคราะห์รูปแบบนี้ เช่น การวิเคราะห์ข้อมูลของลูกค้า การวิเคราะห์ประเมินซัพพลายเออร์ หรือการวิเคราะห์หาสาเหตุของปัญหาการล่าช้าในการจัดส่ง เป็นต้น ในลำดับที่สูงขึ้นถ้าต้องการคาดการณ์หรือพยากรณ์แนวโน้มในอนาคตโดยการใช้ข้อมูลสถิติในอดีต สามารถที่จะใช้วิธีการวิเคราะห์เชิงพยากรณ์ (predictive analytics) ที่ใช้วิธีอนุกรมเวลา (time series method) หรือการเรียนรู้ของเครื่อง การวิเคราะห์ในลักษณะนี้ได้แก่การพยากรณ์อุปสงค์ของลูกค้า สำหรับลำดับสูงสุดคือการวิเคราะห์เพื่อการตัดสินใจ (prescriptive analytics) เพราะนอกจากจะทำการวิเคราะห์ข้อมูลแล้ว ยังจะต้องให้คำแนะนำเกี่ยวกับ กลยุทธ์ แนวทางตัดสินใจ หรือแนวปฏิบัติที่เหมาะสมที่สุด เพื่อปรับปรุงประสิทธิภาพ หรือ เพื่อลดต้นทุนต่าง ๆ ตัวอย่างของการวิเคราะห์เพื่อการตัดสินใจ เช่น การวิเคราะห์ข้อมูลโลจิสติกส์เพื่อจัดเส้นทางการขนส่ง เป็นต้น (Waller & Fawcett, 2013; Souza, 2014)

4. ผลลัพธ์จากการวิเคราะห์สามารถแสดงผลโดยใช้แดชบอร์ด (dashboards) ที่แสดงผลในรูปแบบของการสรุปข้อมูลสำคัญต่าง ๆ ด้วยการใช้เครื่องมือสื่อสารที่ช่วยให้เข้าใจได้ง่าย และสามารถอัปเดตได้อย่างสม่ำเสมอ เช่น การใช้กราฟและแผนภาพเพื่อแสดงแนวโน้มและความสัมพันธ์ของข้อมูล การใช้แผนที่เพื่อแสดงข้อมูลที่เกี่ยวข้องกับตำแหน่งหรือภูมิศาสตร์ เช่น เส้นทางการขนส่ง เป็นต้น

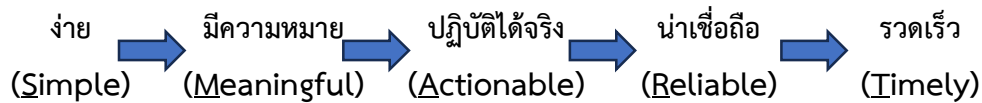
5. ถ้ามีการวิเคราะห์เพื่อการตัดสินใจ ผลลัพธ์จากการวิเคราะห์จะได้เป็นแนวปฏิบัติเพื่อปรับปรุงประสิทธิภาพในการดำเนินการ หลังจากนั้นจะต้องมีการนำแนวปฏิบัตินั้นไปดำเนินการ เช่น การตัดสินใจในการสั่งซื้อหรือการจัดการสต็อกเพื่อให้มีต้นทุนต่ำที่สุด

6. หลังจากนั้นจะต้องมีการติดตามและประเมินผล เริ่มจากการติดตามผลลัพธ์และการดำเนินการให้เป็นไปตามแผนที่วางไว้ โดยจะต้องมีการประเมินผลด้วยการประเมินประสิทธิภาพของการดำเนินการเพื่อหาจุดที่ต้องแก้ไขเพื่อปรับปรุงในอนาคต (Gunasekaran et al., 2001)

โดยภาพรวมการวิเคราะห์ข้อมูลซัพพลายเชนที่มีประสิทธิภาพช่วยให้องค์กรสามารถตัดสินใจได้อย่างมีข้อมูลพื้นฐาน ปรับปรุงกระบวนการ ลดต้นทุน และเพิ่มประสิทธิภาพโดยรวมในซัพพลายเชน

1.5 หลักการวิเคราะห์ข้อมูลซัพพลายเชนอย่างมีประสิทธิภาพ

ดังที่กล่าวไว้แล้วข้างต้นว่าการวิเคราะห์ข้อมูลซัพพลายเชนคือการประยุกต์ใช้ศาสตร์วิทยาการข้อมูลในระดับต่าง ๆ ของซัพพลายเชนเพื่อทำการปรับปรุงประสิทธิภาพโดยรวมของการจัดการซัพพลายเชนเพื่อสนองความต้องการของลูกค้า โดยที่การวิเคราะห์ข้อมูลซัพพลายเชนได้อย่างมีประสิทธิภาพนั้นจำเป็นจะต้องมีการตระหนักถึงหลักการที่เรียกว่า SMART ดังแสดงในภาพที่ 1.4



ภาพที่ 1.4 หลักการ SMART

ที่มา: ดัดแปลงจาก (Liu, 2022)

ความหมายโดยละเอียดของหลักการ SMART สามารถอธิบายได้ต่อไปนี้ เริ่มจาก **S**imple โมเดลการวิเคราะห์ข้อมูลซัพพลายเชนจะต้องง่ายที่สุดที่เป็นไปได้ และสามารถที่จะให้คำตอบที่ง่ายต่อการเข้าใจ **M**eaningful ผลลัพธ์ของการวิเคราะห์ข้อมูลซัพพลายเชนจะต้องมีความหมาย ซึ่งสามารถให้ข้อมูลเชิงลึกที่มีประโยชน์ต่อการตัดสินใจอย่างมีประสิทธิภาพ **A**ctionable ผลลัพธ์ของการวิเคราะห์ข้อมูลซัพพลายเชนจะต้องสามารถปฏิบัติได้จริง เพื่อทำการแก้ไขปัญหาของซัพพลายเชนที่เกิดขึ้นได้ **R**eliable การวิเคราะห์ข้อมูลซัพพลายเชนจะต้องเชื่อถือได้ สามารถให้ผลลัพธ์การแก้ปัญหาที่สม่ำเสมอและน่าเชื่อถือให้กับการจัดการซัพพลายเชนอย่างมีประสิทธิภาพ **T**imely การวิเคราะห์ข้อมูลซัพพลายเชนจำเป็นที่จะต้องทำอย่างรวดเร็วและผลลัพธ์ที่ได้สามารถแก้ปัญหาที่ซัพพลายเชนกำลังเผชิญอยู่ได้อย่างรวดเร็ว

สรุป

การวิเคราะห์ข้อมูลซัพพลายเชนเป็นกระบวนการที่ใช้ข้อมูลเพื่อทำความเข้าใจและปรับปรุงประสิทธิภาพของซัพพลายเชน ซึ่งประกอบด้วยหลายส่วนสำคัญที่ช่วยให้การวิเคราะห์มีประสิทธิภาพ และสามารถตัดสินใจได้อย่างถูกต้อง หัวใจสำคัญของการวิเคราะห์ข้อมูลคือ การประยุกต์ใช้ศาสตร์วิทยาการข้อมูลที่มีส่วนประกอบเริ่มจาก การกำหนดปัญหา คำถามงานวิจัยและวัตถุประสงค์ การรวบรวมและจัดการข้อมูล การวิเคราะห์ข้อมูล การเรียนรู้ของเครื่อง การแสดงผลการวิเคราะห์ และการสื่อสารการนำไปใช้ในการตัดสินใจ โดยทุกส่วนประกอบมีความสำคัญที่เท่าเทียมส่งผลถึงกัน เช่น การรวบรวมและจัดการข้อมูลที่ต้องเหมาะสมส่งผลให้เกิดการเรียนรู้ของเครื่องที่ถูกต้องแม่นยำ การแสดงผลและการสื่อสารถึงผลการวิเคราะห์ที่ง่ายต่อการเข้าใจ ตรงประเด็นอย่างมีประสิทธิภาพ ส่งผลให้ ผู้ใช้งาน ผู้บริหารสามารถนำไปใช้ในการตัดสินใจให้เกิดประโยชน์ได้ เปรียบเสมือนเดียวกับในมุมมองของการจัดการซัพพลายเชน การได้รับข้อมูลที่ต้องแบบเรียลไทม์เพื่อที่จะนำมาแบ่งปันและทำการวิเคราะห์วางแผนได้อย่างรวดเร็ว นั้น ส่งผลอย่างมากต่อการเพิ่มประสิทธิภาพการจัดการซัพพลายเชนโดยรวมสอดคล้องกับหลักการ 7Rs อันได้แก่ (1) right cost (2) right quantity (3) right quality (4) right time (5) right place และ (6) right customer อย่างไรก็ตามในปัจจุบันเมื่อเทคโนโลยีมีความก้าวหน้า เราสามารถที่จะรวบรวมข้อมูลได้ในปริมาณมาก ๆ และมีความหลากหลาย หรือที่เรียกอีกอย่างหนึ่งว่ามีลักษณะเป็นข้อมูลขนาดใหญ่ การใช้เทคโนโลยีการเก็บและประมวลผลข้อมูลแบบเดิมทำให้เกิดความล่าช้าและขาดประสิทธิภาพ จึงเป็นที่มาของการใช้คลาวด์คอมพิวเตอร์ ควบคู่กับวิทยาการข้อมูลเพื่อรับมือกับข้อมูลขนาดใหญ่ที่มีความหลากหลายและรวดเร็วทั้งในด้านการรวบรวม จัดเก็บ และประมวลผลได้อย่างมีประสิทธิภาพแม่นยำอย่างรวดเร็ว

แบบฝึกหัด

1. จงอธิบายถึงส่วนประกอบหลักของวิทยาการข้อมูล
2. การแสดงผลข้อมูล (Data Visualization) มีความจำเป็นอย่างไรและ ถ้าต้องการแสดงความสัมพันธ์ระหว่าง 2 ปัจจัย ควรจะใช้การแสดงผลด้วยกราฟชนิดอะไร
3. อะไรคือคลาวด์คอมพิวติ้ง และจะถูกนำมาประยุกต์ใช้อย่างไรเพื่อเพิ่มระดับความร่วมมือของซัพพลายเชน
4. อธิบายเปรียบเทียบถึงความแตกต่างระหว่างคลาวด์คอมพิวติ้งที่มีโครงสร้างแบบ Laas กับ Paas
5. อธิบายถึงการเรียนรู้ของเครื่องว่าหมายถึงอะไร และการเรียนรู้ของเครื่องมีส่วนช่วยในการเพิ่มประสิทธิภาพของการจัดการซัพพลายเชนอย่างไร
6. ยกตัวอย่างกรณีศึกษาเกี่ยวกับการใช้หลักการ 7Rs ในการจัดการซัพพลายเชนอย่างมีประสิทธิภาพ
7. อธิบายถึงหลักการ SMART ว่ามีความหมายอย่างไร และมีความจำเป็นอย่างไรสำหรับการวิเคราะห์ข้อมูลซัพพลายเชน

เอกสารอ้างอิง

- Areerakulkan, N., & Pongpech, W. (2021). A Dempster-Shafer big data readiness assessment model. In *Proceedings of the 23rd International Conference on Enterprise Information Systems (ICEIS 2021)* (Vol. 2, pp. 581-585). SCITEPRESS. <https://doi.org/10.5220/0010506205810585>.
- Cao, L. (2017). *Data science: A comprehensive overview*. *ACM Computing Surveys*, 50(3), Article 43, 1-42. <https://doi.org/10.1145/3076253>
- Dhar, V. (2013). Data science and prediction. *Communications of the ACM*, 56(12), 64-73. <https://doi.org/10.1145/2500499>.
- Gunasekaran, A., Patel, C., & Tirtiroglu, E. (2001). Performance measures and metrics in a supply chain environment. *International Journal of Operations & Production Management*, 21(1/2), 71-87. <https://doi.org/10.1108/01443570110358468>.
- Han, J., Kamber, M., & Pei, J. (2011). *Data mining: Concepts and techniques* (3rd ed.). Elsevier.
- Hastie, T., Tibshirani, R., & Friedman, J. H. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer.
- Knaflic, C. N. (2015). *Storytelling with data: A data visualization guide for business professionals*. John Wiley & Sons.
- Kochan, C. G., Nowicki, D. R., & Sauser, B. (2018). Impact of cloud-based information sharing on hospital supply chain performance: A system dynamics framework. *International Journal of Production Economics*, 195, 168-185. <https://doi.org/10.1016/j.ijpe.2017.10.008>.
- Kudyba, S. (2014). *Big data, mining, and analytics: Components of strategic decision making* (1st ed.). Auerbach Publications. <https://doi.org/10.1201/b16666>.
- Liu, K. Y. (2022). *Supply Chain Analytics: Concepts, Techniques and Applications*. (1st ed.) Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-92224-5>.
- Mitchell, T. M. (1997). *Machine learning*. McGraw-Hill.
- Munzner, T. (2014). *Visualization analysis and design*. CRC Press.

เอกสารอ้างอิง (ต่อ)

- Neaga, I., Liu, S., Xu, L., Chen, H., & Hao, Y. (2015). Cloud enabled big data business platform for logistics services: A research and development agenda. In *Decision support systems V: Big data analytics for decision making* (pp. 22-33). Springer. https://doi.org/10.1007/978-3-319-18533-0_3.
- Pan, S. J., & Yang, Q. (2010). A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22(10), 1345–1359. <https://doi.org/10.1109/TKDE.2009.191>.
- Provost, F., & Fawcett, T. (2013). *Data science for business: What you need to know about data mining and data-analytic thinking*. O'Reilly Media.
- Simchi-Levi, D., Kaminsky, P., & Simchi-Levi, E. (2004). *Managing the supply chain: The definitive guide for the business professional*. McGraw-Hill.
- Souza, G. C. (2014). Supply chain analytics. *Business Horizons*, 57(5), 595-605. <https://doi.org/10.1016/j.bushor.2014.06.004>.
- Sutton, R. S., & Barto, A. G. (2018). *Reinforcement learning: An introduction* (2nd ed.). MIT Press.
- Udenio, M., Vatamidou, E., Fransoo, J. C., & Dellaert, N. (2017). Behavioral causes of the bullwhip effect: An analysis using linear control theory. *IIE Transactions*, 49(10), 980-1000. <https://doi.org/10.1080/24725854.2017.1325026>.
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77–84. <https://doi.org/10.1111/jbl.12010>
- Xia, B. S., & Gong, P. (2014). Review of business intelligence through data analysis. *Benchmarking*, 21(2), 300-311. <https://doi.org/10.1108/BIJ-08-2012-0050>.
- Zhu, X. (2005). *Semi-supervised learning literature survey* (Technical Report 1530). Department of Computer Sciences, University of Wisconsin–Madison.

บทที่ 2

ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล Data-Driven Supply Chains

หัวข้อ

- 2.1 ข้อมูลและคุณค่าของข้อมูลในการจัดการซัพพลายเชน
- 2.2 ความแตกต่างของซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล
- 2.3 เทคโนโลยีทำให้เกิดข้อมูลขนาดใหญ่ในซัพพลายเชน
- 2.4 คุณลักษณะของข้อมูลขนาดใหญ่
- 2.5 ข้อมูลขนาดใหญ่ในซัพพลายเชน

ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล หมายถึงระบบซัพพลายเชนที่ใช้ข้อมูลในการตัดสินใจและดำเนินการในทุกขั้นตอน ตั้งแต่การจัดหาวัตถุดิบ การผลิต การจัดเก็บ ไปจนถึงการกระจายสินค้า การใช้ข้อมูลที่ถูกต้องและทันสมัยช่วยเพิ่มประสิทธิภาพ ลดต้นทุน และเพิ่มความพึงพอใจของลูกค้าในกระบวนการทั้งหมด ซึ่งในปัจจุบันการเปลี่ยนแปลงเป็นซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูลกลายเป็นเรื่องที่เป็นอย่างยิ่ง เนื่องมาจากปัจจัยสำคัญหลายประการที่ผลักดันให้ธุรกิจต้องปรับตัวเพื่อความอยู่รอดและความได้เปรียบในการแข่งขันในตลาดที่มีความซับซ้อนและเปลี่ยนแปลงอย่างรวดเร็ว โดยปัจจัยสำคัญที่ทำให้ต้องเปลี่ยนเป็นซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล ได้แก่ (1) ความซับซ้อนที่เพิ่มขึ้นของซัพพลายเชน ในปัจจุบันซัพพลายเชนมีความซับซ้อนมากขึ้นเนื่องจากการทำงานร่วมกับซัพพลายเออร์หลายรายในหลายประเทศ การจัดการซัพพลายเชนจึงจำเป็นต้องใช้ข้อมูลในการประสานงานและการวางแผนที่มีประสิทธิภาพเพื่อให้การดำเนินงานเป็นไปอย่างราบรื่น (2) ความต้องการของลูกค้าที่เปลี่ยนแปลงอย่างรวดเร็ว เนื่องมาจากลูกค้าในปัจจุบันมีความคาดหวังสูงและมีความต้องการที่เปลี่ยนแปลงอย่างรวดเร็ว การใช้ข้อมูลในการคาดการณ์และตอบสนองต่อความต้องการของลูกค้าจึงเป็นสิ่งจำเป็น เพื่อรักษาความพึงพอใจและความภักดีของลูกค้า (3) การเพิ่มขึ้นของการแข่งขันในตลาด ทำให้ธุรกิจต้องเผชิญกับการแข่งขันที่เข้มข้นมากขึ้น การขับเคลื่อนซัพพลายเชนด้วยข้อมูลช่วยให้บริษัทสามารถเพิ่มประสิทธิภาพ ลดต้นทุน และตอบสนองต่อการ

แข่งขันได้อย่างรวดเร็ว ทำให้สามารถอยู่รอดและเติบโตได้ในตลาด (4) การเติบโตของเทคโนโลยีดิจิทัล เนื่องมาจากเทคโนโลยีดิจิทัล เช่น การเรียนรู้ของเครื่อง ปัญญาประดิษฐ์ (AI) และ IoT (Internet of Things) กำลังเข้ามามีบทบาทสำคัญในธุรกิจ และชัพพลายเซนที่ขับเคลื่อนด้วยข้อมูลสามารถนำเทคโนโลยีเหล่านี้มาใช้ในการเพิ่มประสิทธิภาพการดำเนินงานได้ (5) การลดความเสี่ยงในชัพพลายเซน เพราะข้อมูลที่ถูกต้องและทันสมัยช่วยในการคาดการณ์และจัดการความเสี่ยงในชัพพลายเซน เช่น การขาดแคลนวัตถุดิบ ปัญหาการขนส่ง หรือการเปลี่ยนแปลงในตลาด การใช้ข้อมูลช่วยให้ธุรกิจสามารถเตรียมตัวและป้องกันปัญหาได้ดีกว่า (6) ความจำเป็นในการเพิ่มความโปร่งใสและความสามารถในการตรวจสอบ เนื่องมาจากชัพพลายเซนที่ขับเคลื่อนด้วยข้อมูลทำให้การดำเนินงานมีความโปร่งใสและสามารถตรวจสอบได้ตลอดเวลา ซึ่งเป็นสิ่งที่ลูกค้าและผู้มีส่วนได้เสียอื่นๆ คาดหวังในปัจจุบัน โดยเฉพาะอย่างยิ่งในกรณีของสินค้าที่ต้องการการตรวจสอบย้อนกลับ (7) ความต้องการในการลดต้นทุนและเพิ่มประสิทธิภาพ โดยการใช้ข้อมูลในการวิเคราะห์และปรับปรุงกระบวนการในชัพพลายเซนสามารถช่วยลดต้นทุนในการดำเนินงานและเพิ่มประสิทธิภาพได้ เช่น การจัดการสต็อกที่มีประสิทธิภาพ การปรับปรุงกระบวนการผลิต และการเลือกเส้นทางการขนส่งที่ดีที่สุด (8) ปัจจัยอันดับสุดท้ายมาจากการตอบสนองต่อกฎระเบียบและมาตรฐานที่เพิ่มขึ้น ทำให้หลายอุตสาหกรรมต้องปฏิบัติตามกฎระเบียบและมาตรฐานที่เข้มงวดมากขึ้น การใช้ข้อมูลในการตรวจสอบและรายงานการดำเนินงานช่วยให้บริษัทสามารถปฏิบัติตามข้อกำหนดเหล่านี้ได้อย่างมีประสิทธิภาพ

โดยสรุปการเปลี่ยนแปลงเป็นชัพพลายเซนที่ขับเคลื่อนด้วยข้อมูล ไม่ใช่เพียงแค่การปรับปรุงกระบวนการเท่านั้น แต่เป็นการตอบสนองต่อปัจจัยภายนอกที่มีผลกระทบต่อการดำเนินธุรกิจในยุคดิจิทัล ธุรกิจที่สามารถปรับตัวและใช้ข้อมูลอย่างมีประสิทธิภาพจะมีความได้เปรียบในการแข่งขัน สามารถลดความเสี่ยง เพิ่มความพึงพอใจของลูกค้า และสร้างความยั่งยืนในระยะยาวได้ (Schoenherr & Speier-Pero, 2015; Kache & Seuring, 2017)

2.1 ข้อมูลและคุณค่าของข้อมูลในการจัดการชัพพลายเซน

"ข้อมูล" คือ สิ่งต่างๆ ที่ถูกเก็บรวบรวมไว้เพื่อสามารถนำไปถูกประมวลผลช่วยในการตัดสินใจ อย่างไรก็ตามปัจจุบันได้มีการให้คำนิยามของ ข้อมูลในยุคปัจจุบันเป็นข้อมูลดิจิทัล (digital data) หมายถึง ข้อมูลที่ถูกสร้างขึ้น รวบรวม และจัดเก็บในรูปแบบดิจิทัล ซึ่งสามารถถูกเข้าถึงและ

ประมวลผลได้โดยใช้เทคโนโลยีดิจิทัล เช่น คอมพิวเตอร์ อินเทอร์เน็ต หรือระบบคลาวด์ เป็นต้น ข้อมูลดิจิทัลมีความสำคัญมากในการทำให้เกิดการวิเคราะห์ นำข้อมูลไปใช้ประโยชน์ และเชื่อมโยงกับกระบวนการธุรกิจและองค์กรต่าง ๆ ในยุคปัจจุบัน

กรณีศึกษาบริษัท Coca-Cola

ในปัจจุบันข้อมูลมีบทบาทสำคัญมากขึ้นในการดำเนินธุรกิจ ตัวอย่างเช่น Coca-Cola เป็นบริษัทขายเครื่องดื่มที่เป็นที่รู้จักกันทั่วโลก ได้ปล่อยโปรแกรมสมัครสมาชิกผ่านระบบดิจิทัล ซึ่งส่งผลให้ Coca-Cola สามารถเชื่อมต่อกับลูกค้าของตนได้อย่างมีประสิทธิภาพช่วยให้สามารถรวบรวมข้อมูลยอดขายของสมาชิกที่สมัครเข้ามา นอกจากนั้นยังทำการรวบรวมความคิดเห็นของลูกค้าผ่านทางแอปพลิเคชันหรือโซเชียลเน็ตเวิร์ก Coca-Cola ใช้ข้อมูลที่ถูกรวบรวมมาในการพัฒนาผลิตภัณฑ์ สร้างโฆษณาที่ปรับแต่งให้เหมาะกับกลุ่มเป้าหมายต่าง ๆ รวมถึงการสร้างแอปพลิเคชันต่างๆ ยกตัวอย่างเช่นแอปพลิเคชันที่ชื่อ “Coke On” ซึ่งถูกพัฒนาขึ้นเพื่อช่วยผู้ขายในการกำหนดตำแหน่งที่ตั้งของตู้ขายน้ำอัดลมหยอดเหรียญให้สอดคล้องกับยอดขายตามความต้องการของลูกค้าได้ กลยุทธ์ด้านข้อมูลนี้ไม่เพียงแต่ช่วยเพิ่มยอดขายเท่านั้นแต่ยังลดต้นทุนและเพิ่มความพึงพอใจให้กับลูกค้าในสถานที่ภายในเวลาที่เหมาะสม

ดังที่ได้กล่าวไว้ในเนื้อหาบางส่วนของบทนำการวิเคราะห์ข้อมูลชัฟฟลายเซน เพื่อให้เกิดการจัดการชัฟฟลายเซนที่มีประสิทธิภาพสอดคล้องกับวัตถุประสงค์ 7Rs องค์กรแต่ละองค์กรจำเป็นต้องบริหารจัดการและประสานงานในกระบวนการจัดทำกับองค์กรอื่นๆในชัฟฟลายเซนอย่างมีประสิทธิภาพ โดยจะต้องมีการแลกเปลี่ยนข้อมูลที่จำเป็นอย่างราบรื่น ทันเวลา และถูกต้องตลอดทั้งกระบวนการจัดทำ เช่น ถ้าปราศจากข้อมูลการพยากรณ์ยอดขายของลูกค้าสุดท้าย (end customers) ที่ถูกส่งมาจากผู้ค้าปลีก ข้อมูลระดับสินค้าคงคลังปัจจุบันของชัฟฟลายเออร์ คำส่งและคำปลีก จะส่งผลให้เป็นเรื่องยากมากต่อผู้ผลิตที่จะตัดสินใจอย่างมีประสิทธิภาพเพื่อวางแผนและกำหนดตารางการผลิตที่เหมาะสม นอกจากนี้ข้อมูลที่ไม่แม่นยำที่ได้รับจากการแลกเปลี่ยนข้อมูล เช่น ความคาดหวังในการตลาดที่ตลาดเคลื่อนเล็กน้อยที่ปลายน้ำอาจทำให้เกิดผลกระทบจากปรากฏการณ์แล้มา ไปยังองค์กรต่างๆทางต้นน้ำของชัฟฟลายเซน (ดูภาพ 2.1) อันจะส่งผลกระทบต่อประสิทธิภาพของทั้งชัฟฟลายเซน และส่งผลให้เกิดปริมาณการผลิตที่เพิ่มขึ้นไม่สอดคล้องกับความต้องการแท้จริงของตลาด ต้นทุนการจัดการสินค้าคงคลังที่สูงตาม กระทบไปยังต้นทุนต่อหน่วยสินค้าที่สูงกว่าแบรนด์

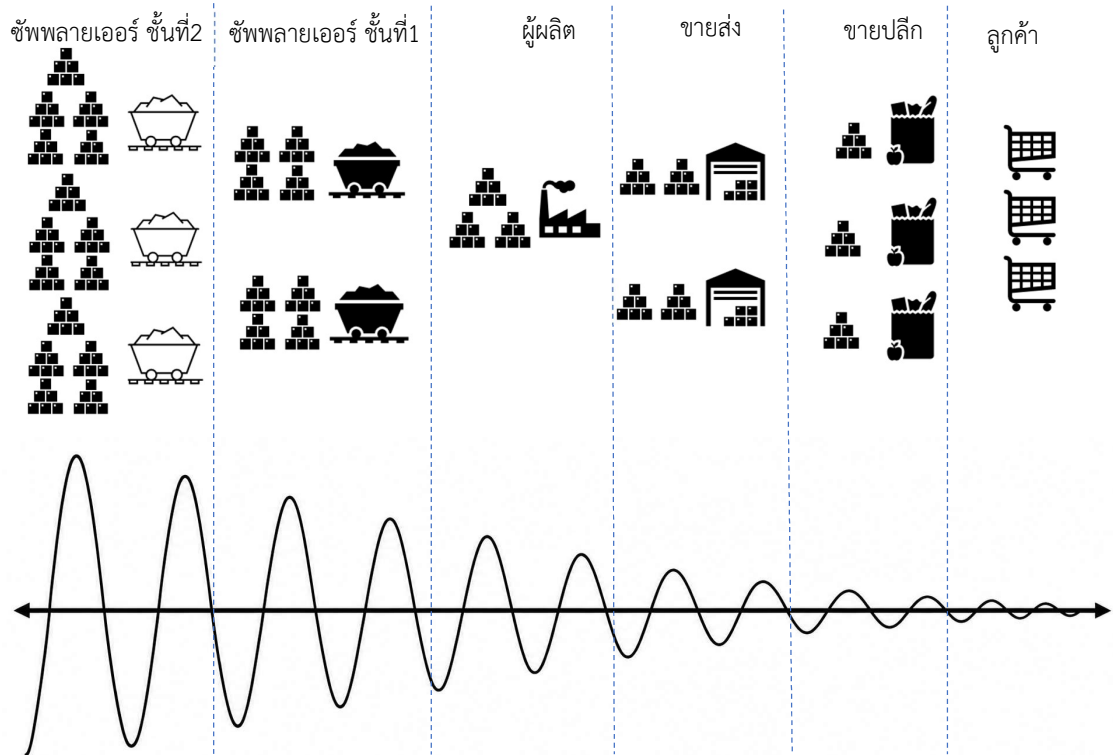
คู่แข่งในในตลาดทำให้สูญเสียความได้เปรียบในการแข่งขัน และไม่สามารถอยู่รอดได้ในระยะยาว (Lee , Padmanabhan, & Whang, 1997)

ด้วยเหตุนี้เองหลาย ๆ องค์กรตระหนักถึงความสำคัญของข้อมูลสำหรับการบริหารจัดการใน ชัฟฟลายเซน ดังนั้นจึงมีการนำโปรแกรมต่าง ๆ มาประยุกต์ใช้เพื่อรวบรวมและแลกเปลี่ยนข้อมูลที่ จำเป็นในชัฟฟลายเซน เช่น ระบบ ERP (การวางแผนทรัพยากรองค์กร) ที่สามารถแยกย่อยออกเป็น โมดูลย่อยตามลักษณะการใช้งานต่างๆได้แก่ ระบบ WMS (การจัดการคลังสินค้า) ระบบ TMS (การ จัดการขนส่ง) และระบบ CRM (การบริหารความสัมพันธ์กับลูกค้า) เป็นต้น โดยที่ผู้รับผิดชอบหรือ ผู้จัดการต้องตัดสินใจอย่างรอบคอบเกี่ยวกับระบบที่จะใช้ให้สอดคล้องกับประเภทรธุรกิจ วัตถุประสงค์ และกระบวนการเดิมที่มีอยู่แล้ว มิฉะนั้นปัญหาที่สำคัญอาจเกิดขึ้นในชัฟฟลายเซน ซึ่งอาจส่งผลให้เกิด การผลิตที่ชะงักเพราะการจัดหาวัสดุหรือวัตถุดิบที่ช้าลงอย่างรุนแรง อันนำไปสู่การสูญเสียลูกค้าใน ท้ายที่สุด

อย่างไรก็ตามการที่จะพึ่งเทคโนโลยีโดยการใช้ซอฟต์แวร์ ERP เพียงอย่างเดียว อาจไม่ใช่ แนวทางแก้ปัญหาก็ถูกต้องเสมอไป เพราะยังมีปัจจัยอื่น ๆ ที่ทำให้การแบ่งปันข้อมูลในชัฟฟลายเซน อย่างมีประสิทธิภาพไม่เกิดขึ้น เช่น การขาดความเชื่อใจกันระหว่างสมาชิกในชัฟฟลายเซนส่งผลให้ เกิดความไม่เต็มใจในการเปิดเผยข้อมูลที่อ่อนไหว (เช่น ข้อมูลการขายและบัญชีลูกค้า) หรือ ระบบที่ ไม่สอดคล้องกันในชัฟฟลายเซน (เช่น การใช้ระบบปฏิบัติการ (operating system) ที่แตกต่างกันโดย ชัฟฟลายเออร์) ทำให้เกิดอุปสรรคในการเขียนโปรแกรม เป็นต้น

ดังนั้น เพื่อที่จะเอาชนะปัญหาต่างๆเหล่านี้เพื่อการแบ่งปันข้อมูลที่มีประสิทธิภาพ องค์กรต้อง สร้าง 2 สิ่งที่สำคัญ สิ่งแรกคือการสร้างพันธมิตรกับสมาชิกอื่น ๆ ในชัฟฟลายเซน ทำให้เกิดความเชื่อ ใจและพร้อมที่จะแบ่งปันข้อมูลที่สำคัญและทันเวลากับองค์กรของตนเองรวมทั้งกับสมาชิกอื่น ๆ แน่แน่นอนว่าการสร้างพันธมิตรและความเชื่อใจนั้นอาจใช้เวลานานมากแต่ผลที่จะได้รับนั้นคุ้มค่ากับ ความพยายาม สิ่งที่สองคือ การสร้างพันธมิตรและการประสานงานของหน่วยงานต่าง ๆ ภายใน องค์กรเอง ทำให้กระบวนการแบ่งปันข้อมูลภายในนั้นง่ายขึ้น โดย ปรับปรุง ส่งเสริม การทำงานร่วม ข้ามสายงาน (cross-functional collaboration) ถึงแม้ว่าบางครั้งอาจต้องมีการเปลี่ยนแปลง โครงสร้างองค์กรเดิม หรือการจัดการกระบวนการธุรกิจใหม่ (reengineering) อีกครั้ง หรือแม้แต่การ รื้อวัฒนธรรมโครงสร้างองค์กรเดิม ผ่านการสนับสนุนและความมุ่งมั่นของผู้บริหารระดับสูง แต่ก็เป็น สิ่งจำเป็นที่จะทำให้เกิดการเปลี่ยนแปลงจากชัฟฟลายเซนแบบดั้งเดิม ก้าวสู่การเป็นชัฟฟลายเซนที่

ขับเคลื่อนด้วยข้อมูล ที่มีการแบ่งปันข้อมูลให้เป็นไปอย่างมีประสิทธิภาพและราบรื่นทั้งภายในและภายนอกขององค์กร



ภาพที่ 2.1 การขยายผลกระทบของปรากฏการณ์แล้มาเกิดจากความต้องการสินค้าของลูกค้าเพียงเล็กน้อย แต่ส่งผลกระทบต่อปริมาณสั่งซื้อและสินค้าคงคลังที่เพิ่มขึ้นอย่างมากไม่ตรงกับความต้องการที่แท้จริง เมื่อคำสั่งซื้อเคลื่อนจากปลายน้ำ (ลูกค้า) ไปยังต้นน้ำ (ซัพพลายเออร์)

ที่มา: ดัดแปลงจาก (Liu, 2022)

2.2 ความแตกต่างของซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล

ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล และซัพพลายเชนแบบดั้งเดิมที่ยังไม่เห็นความสำคัญของการนำข้อมูลมาใช้วิเคราะห์และช่วยในการตัดสินใจ มีความแตกต่างกันอย่างมีนัยสำคัญในหลายด้าน โดยเฉพาะในการใช้ข้อมูลเพื่อการตัดสินใจและการดำเนินการ ซึ่งสามารถแบ่งออกเป็นประเด็นหลัก ๆ แสดงดังตารางที่ 2.1 ดังนี้:

ตารางที่ 2.1 ตารางเปรียบเทียบแนวทางการดำเนินการ

ประเด็นข้อแตกต่าง	แนวทางการดำเนินการ	
	ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล	ซัพพลายเชนแบบดั้งเดิม
1. การใช้ข้อมูลในการตัดสินใจ	การตัดสินใจในทุกขั้นตอนของซัพพลายเชนจะอิงจากข้อมูลที่ถูกต้องและเป็นปัจจุบัน ข้อมูลเหล่านี้มาจากแหล่งต่าง ๆ เช่น ข้อมูลการขาย การผลิต สต็อกสินค้า และข้อมูลจากซัพพลายเออร์ ใช้การวิเคราะห์ข้อมูล (Data Analytics) เพื่อคาดการณ์แนวโน้มในอนาคตและให้คำแนะนำที่เป็นประโยชน์ในการตัดสินใจ เช่น การคาดการณ์ความต้องการของตลาดหรือการจัดการสต็อก	การตัดสินใจอาจอิงจากประสบการณ์ สัญชาตญาณ หรือข้อมูลที่ไม่ได้ถูกวิเคราะห์อย่างละเอียด ทำให้การตัดสินใจอาจไม่แม่นยำและเกิดความเสี่ยงสูง ไม่มีการใช้เทคโนโลยีขั้นสูงในการวิเคราะห์ข้อมูล ทำให้ขาดความสามารถในการคาดการณ์แนวโน้มในอนาคต
2. การตอบสนองต่อการเปลี่ยนแปลง	สามารถตอบสนองต่อการเปลี่ยนแปลงของตลาดหรือความต้องการของลูกค้าได้อย่างรวดเร็ว เพราะมีข้อมูลที่อัปเดตและสามารถวิเคราะห์ได้แบบเรียลไทม์ สามารถปรับกลยุทธ์หรือกระบวนการต่าง ๆ ได้อย่างทันทีตามข้อมูลที่ได้รับ	มักจะมีการตอบสนองที่ช้าต่อการเปลี่ยนแปลง เพราะข้อมูลอาจล่าช้าหรือไม่ครบถ้วน ทำให้ไม่สามารถทำการปรับเปลี่ยนได้ทันทั่วทั้งการเปลี่ยนแปลงมักต้องใช้เวลาอันเนื่องมาจากขาดข้อมูลที่สนับสนุนการตัดสินใจ

ตารางที่ 2.1 ต่อ

ประเด็นข้อแตกต่าง	แนวทางการดำเนินการ	
	ชัฟฟลายเซนที่ขับเคลื่อนด้วยข้อมูล	ชัฟฟลายเซนแบบดั้งเดิม
3. ประสิทธิภาพในการดำเนินการ	มีประสิทธิภาพสูงในการดำเนินการ เนื่องจากใช้ข้อมูลเพื่อเพิ่มประสิทธิภาพในทุกขั้นตอน เช่น การลดเวลานำสินค้าเข้า, การลดต้นทุนในการจัดการสต็อก, และการเพิ่มความแม่นยำในการจัดส่ง สามารถตรวจสอบและปรับปรุงกระบวนการอย่างต่อเนื่องด้วยข้อมูลที่ถูกระบุรวบรวมและวิเคราะห์	อาจประสบกับความไม่มีประสิทธิภาพในกระบวนการต่าง ๆ เนื่องจากขาดข้อมูลที่ทันสมัยและการวิเคราะห์ที่แม่นยำ การปรับปรุงกระบวนการทำได้ยาก เนื่องจากขาดข้อมูลในการวัดผลและการวิเคราะห์
4. ความโปร่งใสและการสื่อสาร	มีความโปร่งใสสูง เนื่องจากสามารถแบ่งปันข้อมูลระหว่างฝ่ายต่าง ๆ ในชัฟฟลายเซนได้อย่างรวดเร็วและแม่นยำ การสื่อสารภายในองค์กรและกับชัฟฟลายเออร์เป็นไปอย่างมีประสิทธิภาพ เนื่องจากทุกฝ่ายสามารถเข้าถึงข้อมูลที่จำเป็นได้	อาจขาดความโปร่งใสเนื่องจากข้อมูลอาจกระจายไม่ทั่วถึงหรือเข้าถึงได้ยาก ทำให้เกิดการสื่อสารที่ไม่ครบถ้วน การสื่อสารระหว่างฝ่ายต่าง ๆ อาจเกิดความล่าช้าหรือเกิดความเข้าใจผิด เนื่องจากข้อมูลไม่ครบถ้วนหรือไม่ทันสมัย

ตารางที่ 2.1 ต่อ

ประเด็นข้อแตกต่าง	แนวทางการดำเนินการ	
	ชัฟฟลายเซนที่ขับเคลื่อนด้วยข้อมูล	ชัฟฟลายเซนแบบดั้งเดิม
5. ความเสี่ยงและการจัดการปัญหา	สามารถลดความเสี่ยงได้ดี เพราะสามารถคาดการณ์และเตรียมการรับมือกับปัญหาที่อาจเกิดขึ้นได้ล่วงหน้า เช่น การขาดแคลนสินค้า หรือการล่าช้าในการจัดส่ง มีการใช้ข้อมูลในการวิเคราะห์ความเสี่ยงและหาทางแก้ไข ปัญหาอย่างมีประสิทธิภาพ	มีความเสี่ยงสูงกว่าในการจัดการปัญหา เพราะการตัดสินใจอาจไม่ได้อิงจากข้อมูลที่แม่นยำ ทำให้เกิดความผิดพลาดได้ง่าย การจัดการปัญหามักเป็นแบบปฏิกิริยา (Reactive) มากกว่าการป้องกัน (Proactive) เนื่องจากขาดข้อมูลในการวิเคราะห์ล่วงหน้า

ที่มา: จัดทำโดยผู้แต่ง

โดยสรุป ชัฟฟลายเซนที่ขับเคลื่อนด้วยข้อมูล จะมีความทันสมัย มีความสามารถในการปรับตัวและตอบสนองต่อความต้องการของตลาดได้ดี และมีประสิทธิภาพในการดำเนินการสูงกว่าชัฟฟลายเซนธรรมดา ที่มีกึ่งพาประสบการณ์และการคาดเดามากกว่า ซึ่งทำให้เกิดความเสี่ยงและการดำเนินการที่ไม่แน่นอนมากขึ้น

2.3 เทคโนโลยีที่ทำให้เกิดข้อมูลขนาดใหญ่ในชัฟฟลายเซน

ในชัฟฟลายเซนมีแหล่งข้อมูลมากมายที่มีอยู่ไม่เพียงแค่ภายในองค์กรของตัวเอง แต่ยังรวมถึงข้ามขอบเขตของบริษัทด้วย การแลกเปลี่ยนข้อมูลระหว่างสมาชิกในชัฟฟลายเซนเป็นสิ่งสำคัญสำหรับการประสานงานและการทำงานร่วมกันอย่างมีประสิทธิภาพ หากสมาชิกในแต่ละชั้น (tier) ภายในชัฟฟลายเซนแบ่งปันข้อมูลอย่างเหมาะสมกับสมาชิกชั้นอื่น ๆ ก็จะสามารถให้การตอบสนองที่รวดเร็วและถูกต้องได้ เช่น เจ้าหน้าที่ขายรถจักรยานยนต์แบ่งปันข้อมูลยอดขายในเวลาที่เหมาะสมกับผู้ผลิตรถจักรยานยนต์ ผู้ผลิตรถจักรยานยนต์จึงใช้ข้อมูลการขายเพื่อปรับแผนการผลิตที่โรงงานของตน

แผนการผลิตที่สำคัญจะต้องถูกสื่อสารอย่างรวดเร็วกับซัพพลายเออร์เช่นเดียวกัน ซึ่งจากนั้นพวกเขาจึงพยายามปรับแผนการผลิตของตนเองเพื่อส่งมอบวัตถุดิบแบบทันเวลาพอดี (JIT) ไปยังโรงงานผู้ผลิตรถจักรยานยนต์

การนำเทคโนโลยีสารสนเทศ (IT) มาใช้ ได้เป็นที่ยอมรับว่าช่วยลดกระบวนการในซัพพลายเชน ช่วยปรับปรุงการมองเห็น ช่วยให้เกิดการแบ่งปันข้อมูลที่ดี และส่งเสริมให้การตัดสินใจเป็นไปอย่างมีประสิทธิภาพมากขึ้น เทคโนโลยีที่เปลี่ยนแปลงอย่างรวดเร็วโดยเฉพาะเทคโนโลยีอัจฉริยะ (smart technology) ได้ทำให้เกิดข้อมูลที่มีขนาดใหญ่มาก ซึ่งจำเป็นต้องถูกจัดการตามขั้นตอนของการวิเคราะห์ข้อมูลที่กล่าวไว้ในเบื้องต้นได้แก่ การดึงข้อมูล รวบรวมข้อมูล ทำความสะอาดข้อมูล จัดเก็บข้อมูล การวิเคราะห์ข้อมูล จนไปถึงการนำผลการวิเคราะห์ไปใช้วางแผนตัดสินใจ เทคโนโลยีเหล่านี้จึงถือว่าเป็นต้นกำเนิดของข้อมูลที่เป็นในการวิเคราะห์เพื่อเพิ่มประสิทธิภาพลดต้นทุนในการจัดการซัพพลายเชน เทคโนโลยีเหล่านี้ เช่น

1. การแลกเปลี่ยนข้อมูลทางอิเล็กทรอนิกส์ (Electronic Data Interchange (EDI))

คือระบบการแลกเปลี่ยนข้อมูลทางอิเล็กทรอนิกส์ เป็นระบบที่ใช้แทนใบสั่งซื้อที่ใช้กระดาษ ใบกำกับสินค้า การส่งจดหมายทางไปรษณีย์ แฟกซ์ และอีเมลด้วยรูปแบบอิเล็กทรอนิกส์มาตรฐาน เพื่อแลกเปลี่ยนเอกสารธุรกิจผ่านเครือข่ายคอมพิวเตอร์ไปยังองค์กรอื่น ๆ อย่างทันที ระบบ EDI ช่วยให้ประหยัดเวลาและลดข้อผิดพลาดที่เกิดขึ้นจากการจัดการโดยมนุษย์ ตัวอย่างของเอกสารที่แลกเปลี่ยนผ่าน EDI รวมถึงใบสั่งซื้อ ใบกำกับสินค้าและเอกสารสถานะการจัดส่ง ข้อมูลสต็อก และใบขนส่งสินค้า (Chen, Li, & Wang, 2020)

2. ระบบบาร์โค้ด (Barcoding System) เป็นระบบที่ประกอบด้วยฮาร์ดแวร์บาร์โค้ด (เช่น

สแกนเนอร์แบบพกพา เครื่องคอมพิวเตอร์โมบายล์ และเครื่องพิมพ์) ใช้เพื่อสแกนและรับรู้ข้อมูลสินค้า ผ่านซอฟต์แวร์ที่สนับสนุนในการอ่าน ถอดรหัส แปลความ และส่งข้อมูลไปยังฐานข้อมูลส่วนกลาง ในปัจจุบันระบบบาร์โค้ดได้รับการนำมาใช้อย่างแพร่หลายในหลายสาขาธุรกิจ เช่น เมื่อชำระสินค้าที่ซูเปอร์มาร์เก็ต หรือการใช้สแกนเนอร์รหัส QR-Code บนโทรศัพท์มือถือเพื่อตรวจสอบราคาสินค้าและทำการซื้อสินค้า การใช้ระบบบาร์โค้ดมีประโยชน์ตรงที่สามารถเก็บข้อมูลโดยอัตโนมัติให้ข้อมูลแบบเรียลไทม์และแม่นยำมาก และลดความเสี่ยงของข้อผิดพลาดจากมนุษย์ลงได้ เราสามารถใช้ระบบบาร์โค้ดในซัพพลายเชนเพื่อจัดการสินค้าขาเข้าและขาออกในคลังสินค้า บันทึกกระตือรือร้นสินค้าคงคลังโดยอัตโนมัติ ติดตามรายละเอียดการขนส่ง สถานะในการขนส่ง และการกระจายสินค้า

หรือการบริการ (Manthou & Vlachopoulou, 2001; McFarlane & Sheffi, 2003; Razak et al., 2023; Singh et al., 2023)

3. การระบุเอกลักษณ์ด้วยคลื่นวิทยุ (Radio-frequency Identification (RFID)) เป็นระบบระบุตัวด้วยคลื่นวิทยุ เป็นระบบที่ใช้คลื่นวิทยุเพื่อถ่ายโอนข้อมูลโดยอัตโนมัติ ทำให้ผู้ใช้สามารถระบุตำแหน่งสินค้าและทรัพย์สินได้อย่างรวดเร็ว เทคโนโลยีนี้ใช้แท็ก (RFID Tags) หรือ smart labels ซึ่งข้อมูลและข้อมูลสารสนเทศถูกบันทึกและเข้ารหัสไว้ RFID Tags จะถูกติดตั้งบนสิ่งของเพื่อติดตามได้โดยใช้เครื่องอ่านและเสาอากาศ RFID Tags มีทั้งแบบที่มีแบตเตอรี่ และไม่มีแบตเตอรี่ โดยจะรับพลังงานจากคลื่นวิทยุที่สร้างขึ้นโดยเครื่องอ่าน ข้อมูลสามารถถูกส่งผ่านคลื่นวิทยุไปยังเครื่องอ่าน/เสาอากาศ แล้วถูกถ่ายโอนไปยังระบบคอมพิวเตอร์หลักและจัดเก็บในฐานข้อมูลเพื่อการวิเคราะห์เพิ่มเติม ต่างจากระบบบาร์โค้ด ข้อมูลจาก RFID Tags สามารถอ่านได้ในแนวที่ไม่เป็นเส้นตรง จึงทำให้มีประโยชน์มากในการประยุกต์ใช้ในหลายๆ งาน เช่น การติดตามสินค้าในคลังสินค้า การจัดการวัสดุและสินค้าในการขนส่ง การติดตามยานพาหนะ การเก็บเงินผ่านระบบทางด่วน และระบบระบุตำแหน่งแบบเรียลไทม์ (Chanchaichujit et al., 2020; Tan & Sidhu, 2022; Unhelkar et al., 2022)

4. ระบบวางแผนทรัพยากรองค์กร (Enterprise Resource Planning (ERP)) ระบบการจัดการทรัพยากรองค์กรรวมกระบวนการธุรกิจที่แตกต่างกันและขั้นตอนหลายขั้นตอนของซัพพลายเชนด้วยฐานข้อมูลร่วมเพื่อให้การสื่อสารข้อมูลระหว่างกันเป็นไปอย่างราบรื่นเรียลไทม์และแม่นยำ ระบบ ERP ประกอบด้วยโมดูล ERP ต่าง ๆ เช่น การเงินและบัญชี การจัดการทรัพยากรมนุษย์ (HRM) การขายและการตลาด การจัดการสต็อก การบริหารความสัมพันธ์กับลูกค้า (CRM) และการบริหารซัพพลายเชน (SCM) ERP โดยทั่วไปใช้โครงสร้างข้อมูลที่กำหนดไว้เพื่อสร้าง เก็บ และแบ่งปันข้อมูลระหว่างกระบวนการหลัก ๆ เหล่านี้ ซึ่งช่วยให้มีการรายงานอย่างรวดเร็ว กำจัดความซ้ำซ้อนของข้อมูลและข้อผิดพลาด และเพิ่มความคงสภาพของข้อมูล ในปัจจุบัน ระบบ ERP ได้รับการพัฒนาเพื่อรองรับเทคโนโลยีที่เกิดขึ้นอย่างต่อเนื่อง เช่น คลาวด์คอมพิวติ้ง (cloud ERP) และสมาร์ทโฟน (mobile ERP) ที่มอบคุณการเข้าถึงระบบ ERP ได้สะดวกยิ่งขึ้น มีค่าใช้จ่ายต่ำกว่า และสามารถเข้าถึงระบบได้จากที่ใดก็ตามที่บริษัทผู้ขาย ERP ขึ้นนำบนตลาดรวมถึง SAP, Oracle, Sage, Microsoft และ Infor (Tarigan et al., 2021; Qureshi, 2022)

5. ระบบสารสนเทศภูมิศาสตร์ (Geographic Information System (GIS)) ระบบสารสนเทศทางภูมิศาสตร์เป็นระบบที่ใช้ในการจับข้อมูลทางภูมิศาสตร์ จัดเก็บ และวิเคราะห์เพื่อการแก้ปัญหาและการตัดสินใจที่มีประสิทธิภาพเกี่ยวกับการวิเคราะห์ทางพื้นที่ GIS สามารถใช้ในชีพพลายเซนเพื่อติดตามจัดการทรัพยากร รวมไปถึงการตัดสินใจในการเลือกสถานที่ได้ เช่น ผู้ทำงานในชีพพลายเซนสามารถใช้ GIS เพื่อกำหนดสถานที่สำหรับอาคาร ติดตามรถบรรทุกหรือเรือ และมองเห็นภาพของสินค้าว่าอยู่ที่ไหน ในปัจจุบัน ระบบ GIS ได้เริ่มเปลี่ยนแปลงเป็นแอปพลิเคชันสถานที่ฉลาด (location intelligence, LI) ที่รวมแหล่งข้อมูลที่มีการเพิ่มความคุ้มค่า (เช่น ข้อมูลทางสถิติประชากร ข้อมูลอากาศ และการสตรีมข้อมูลแบบเรียลไทม์) และรองรับข้อมูลสถานที่ซึ่งละเอียดยิ่งขึ้นเพื่อใช้ในการวิเคราะห์สำหรับการปรับปรุงและการทำนาย เช่น LI สามารถใช้โดยผู้จัดการโลจิสติกส์เพื่อปรับปรุงเส้นทางการจัดส่งโดยขึ้นอยู่กับข้อมูลการจราจรแบบเรียลไทม์และตำแหน่งของลูกค้า นอกจากนี้ ระบบ GIS ยังถูกใช้เพื่อเพิ่มความยืดหยุ่นของชีพพลายเซน โดยเฉพาะอย่างยิ่งเมื่อเผชิญกับการหยุดชะงัก เช่น ภัยพิบัติทางธรรมชาติ มีการนำเสนอกรณีศึกษาหลายกรณีที่ใช้ GIS เพื่อทำแผนที่ความเปราะบางและเพิ่มประสิทธิภาพในการฟื้นฟูจากเหตุการณ์เหล่านี้ (Rodríguez-Espíndola et al., 2018; Koot et al., 2021)

6. ระบบไซเบอร์-กายภาพ (Cyber Physical System (CPS)) เป็นระบบที่เชื่อมโยงโลกทางกายภาพกับโลกเสมือนโดยใช้คอมพิวเตอร์และเครือข่ายที่ฝังตัว ในระบบ CPS สิ่งที่เป็นกายภาพ/กระบวนการจะถูกตรวจสอบ ควบคุม และปรับปรุงผ่านอัลกอริทึมคอมพิวเตอร์ ที่ผูกติดกับเซนเซอร์ แอ็กทูเอเตอร์ เครือข่ายสื่อสาร ซอฟต์แวร์ และคอมพิวเตอร์ เปรียบเสมือนอินเทอร์เน็ตเปลี่ยนแปลงวิธีการที่เรามีการติดต่อกัน ระบบ CPS เปลี่ยนแปลงวิธีการที่โลกทางกายภาพและโลกเสมือนมีปฏิสัมพันธ์กัน การประยุกต์ใช้งานของ CPS ที่สำคัญรวมถึงด้านพลังงาน (เช่น ระบบกริดอัจฉริยะ) ด้านสุขภาพ (เช่น การตรวจสุขภาพทางการแพทย์) การเคลื่อนไหว (เช่น ระบบรถยนต์อัจฉริยะ) และการผลิต (เช่น ระบบควบคุมกระบวนการผลิตอัตโนมัติและ ระบบหุ่นยนต์) CPS สามารถนำมาใช้ในชีพพลายเซนเพื่อส่งเสริมการผลิตที่ฉลาดและมีประสิทธิภาพ การติดตามและตรวจสอบการส่งมอบ การจัดการโลจิสติกส์และควบคุมสินค้าคงคลังให้มีประสิทธิภาพมากขึ้น (Feng & Ye, 2021)

7. อินเทอร์เน็ตของสรรพสิ่ง (Internet of Things (IoT)) คือเครือข่ายของวัตถุทางกายภาพที่สามารถระบุได้โดยเฉพาะ (เช่น เครื่องใช้ในบ้าน เครื่องจักรกลและดิจิทัล ยานพาหนะ อุปกรณ์ ฯลฯ) ที่ฝังเอาไว้ด้วยเซนเซอร์ แอ็กทูเอเตอร์ ซอฟต์แวร์ และอุปกรณ์คอมพิวเตอร์ ทำให้สามารถแลกเปลี่ยนข้อมูลผ่านอินเทอร์เน็ตได้ CPS และ IoT มีความทับซ้อนกันอย่างมีนัยสำคัญ อย่างไรก็ตาม CPS เน้นไปที่การเชื่อมโยงระหว่างกระบวนการคำนวณและโลกทางกายภาพมากกว่า ในขณะที่ IoT เน้นไปที่อุปกรณ์ที่เชื่อมต่อกับอินเทอร์เน็ตและระบบฝังตัวมากกว่า หากเราพิจารณาถึง CPS ในฐานะเป็นขั้นตอนแรกของการรวมดิจิทัลในแนวคิดโดยการเชื่อมโยงระหว่างสิ่งมีชีวิตกับโลกเสมือน IoT จะสร้างมาสู่อนาคตที่ทุกสิ่งเชื่อมต่อกันผ่านอินเทอร์เน็ต ทำให้สามารถเก็บข้อมูลเกี่ยวกับโลกทางกายภาพได้จากทุกที่ และแบ่งปันกับระบบและอุปกรณ์อื่น ๆ หากเราขยายแนวความคิดของ IoT ไปสู่การผลิต IoT ก็จะกลายเป็น Internet of Things ในอุตสาหกรรม (IIoT) ซึ่งยังรู้จักกันในนามของระบบการผลิต Cyber Physical Production System (CPPS) อินเทอร์เน็ตอุตสาหกรรม หรือโรงงานอัจฉริยะ ในทางเดียวกัน หากเราขยายแนวความคิดไปสู่มหาวิทยาลัย IoT ก็จะกลายเป็นบ้านอัจฉริยะด้วยการเชื่อมต่อเครื่องใช้ในบ้าน (เช่น โทรศัพท์ เครื่องปรับอากาศ หอกระจายไฟฟ้า กล้องวงจรปิด เป็นต้น) เมื่ออุปกรณ์หรือเครื่องจักรอัจฉริยะถูกเชื่อมต่อกัน จะสร้างข้อมูล IoT อย่างมากมายซึ่งสามารถนำไปวิเคราะห์ ใช้ประโยชน์ และดำเนินการในเวลาจริงโดยไม่ต้องมีคนเข้ามาเกี่ยวข้อง ในสภาพแวดล้อมของ IoT และ CPS ที่ผสมผสานกับปัญญาประดิษฐ์ การเรียนรู้ของเครื่อง และคลาวด์ เครื่องจักรสามารถทำงานร่วมกับเครื่องจักรอื่น ๆ และมนุษย์ได้ ซึ่งเครื่องจักรสามารถเรียนรู้เกี่ยวกับมนุษย์และปรับตัวตามความต้องการของมนุษย์ได้ มีคำพูดว่าการเกิดของ IoT และ CPS กำลังขับเคลื่อนการเปลี่ยนแปลงที่สำคัญที่สุดในธุรกิจและเทคโนโลยีตั้งแต่สงครามโลกครั้งที่สอง ซึ่งได้สร้างพื้นฐานสำหรับการปฏิวัติอุตสาหกรรมครั้งที่สี่ หรือ Industry 4.0 และกำลังจะก้าวเข้าสู่ครั้งที่ 5 Industry 5.0 ในอนาคตอันใกล้ (Koot et al., 2021)

8. บล็อกเชน (Blockchain) เป็นเทคโนโลยีในการเก็บบันทึกที่อยู่เบื้องหลังของสกุลเงินดิจิทัลบิตคอยน์ ถูกพัฒนาโดยบุคคลที่ไม่เปิดเผยตัวตนที่ใช้นามแฝงว่า Satoshi Nakamoto ในปี ค.ศ. 2008 ต่างจากฐานข้อมูลแบบดั้งเดิม บล็อกเชนเป็นฐานข้อมูลชนิดพิเศษที่เก็บข้อมูลในรูปแบบของบล็อก เมื่อบล็อกใหม่ได้รับข้อมูลใหม่ เข้าสู่การเชื่อมต่อกับบล็อกก่อนหน้านี้ตามลำดับเวลา หนึ่งในคุณสมบัติที่เป็นเอกลักษณ์และข้อดีของบล็อกเชน คือมันเป็นสมุดบัญชีแบบแยกประเภทแบบกระจายของธุรกรรมทั้งหมดข้ามเครือข่ายจำนวนมาก โดยแต่ละโหนดในระบบแบบกระจายมีสำเนา

ของบล็อกเชน ดังนั้นข้อมูลที่ป้อนเข้าไปเป็นข้อมูลที่ไม่สามารถทำให้เสื่อมสลายและไม่สามารถย้อนกลับได้ บล็อกเชนมีการใช้งานที่มีศักยภาพมากมายที่เกินกว่าแคบิตคอยน์และสกุลเงินดิจิทัล บริษัทหลาย ๆ แห่งได้นำเทคโนโลยีนี้มาใช้แล้ว เช่น Walmart Pfizer AIG Siemens Unilever และบริษัทอื่น ๆ อีกมากมาย ตัวอย่างเช่น IBM ได้สร้าง food trust blockchain เพื่อติดตามเส้นทางที่ผลิตภัณฑ์อาหารเดินทางไปถึงสถานที่ต่าง ๆ ในชีพลายเซนอาหารจากต้นน้ำ (เกษตรกร) ไปยังปลายน้ำจนถึงมือผู้บริโภค (Kamath, 2018; Dutta et al., 2020)

2.4 คุณลักษณะของข้อมูลขนาดใหญ่ (Big Data)

การใช้ระบบเทคโนโลยีสารสนเทศและเทคโนโลยีอัจฉริยะในชีพลายเซน พร้อมกับอีคอมเมิร์ซ โซเชียลมีเดีย และอื่น ๆ ทำให้มีปริมาณข้อมูลที่มหาศาลให้กับธุรกิจได้ใช้ประโยชน์ ลักษณะของปริมาณข้อมูลที่มีมากมาย หลากหลาย และรวดเร็ว ทำให้เกิดคำศัพท์ยอดนิยมอย่าง Big Data ในอดีตองค์กรมักจะสามารถในการเก็บข้อมูลปริมาณมากอยู่ในระดับที่จำกัด เนื่องจากข้อจำกัดในการเก็บข้อมูลด้วยวิธีการแบบดั้งเดิมที่มีข้อจำกัดในการประมวลผลข้อมูลที่มีปริมาณมากมาย มีความซับซ้อน และเปลี่ยนแปลงได้อย่างรวดเร็ว ซึ่งในปัจจุบันด้วยพลังการคำนวณที่มากขึ้นและการเก็บข้อมูลด้วยต้นทุนที่ถูกกว่าบนแพลตฟอร์มเช่น data lakes (ระบบฐานข้อมูลที่เก็บข้อมูลได้หลากหลายรูปแบบ ตั้งแต่ ข้อมูลมีโครงสร้าง กึ่งโครงสร้าง ไปจนถึงไม่มีโครงสร้าง) และ Hadoop (ระบบจัดเก็บข้อมูลที่กระจายไฟล์ไปยังเครื่องคอมพิวเตอร์หลายเครื่องซึ่งทำงานร่วมกัน ทำให้เก็บข้อมูลได้ปริมาณมหาศาล) ส่งผลให้ธุรกิจสามารถเข้าถึง เก็บรักษา และประมวลผลข้อมูลได้มากขึ้นในเวลาจริงหรือใกล้เคียงกับเวลาจริงเหมาะสำหรับการตัดสินใจที่มีประสิทธิภาพทันต่อการเปลี่ยนแปลงอย่างรวดเร็วในสถานการณ์ปัจจุบัน (Areerakulkan & Pongpech, 2021)

คุณลักษณะหลักของข้อมูลขนาดใหญ่ (5V's ของ Big Data) คุณลักษณะหลักของข้อมูลขนาดใหญ่ประกอบด้วย 5 คุณลักษณะได้แก่

1. **Volume (ปริมาณ)** ข้อมูลขนาดใหญ่มีปริมาณมหาศาล ซึ่งอาจมีขนาดตั้งแต่เทราไบต์ (TB) ไปจนถึงเซตาไบต์ (ZB) ขึ้นอยู่กับแหล่งข้อมูล เช่น ข้อมูลจากโซเชียลมีเดีย ข้อมูลจากเซ็นเซอร์ ข้อมูลธุรกรรม หรือข้อมูลการคลิกบนเว็บไซต์ เป็นต้น

2. Variety (ความหลากหลาย) ข้อมูลขนาดใหญ่มาในรูปแบบต่าง ๆ ซึ่งสามารถแยกประเภทเป็นสามประเภทหลักได้ดังนี้

2.1 ข้อมูลที่มีโครงสร้าง (Structured Data) หมายถึงข้อมูลที่มีรูปแบบที่แน่นอนหรือเรียงลำดับไว้ในลักษณะที่สามารถจัดเก็บ ค้นหา และประมวลผลได้อย่างราบรื่น และไม่มีปัญหา เช่นเอกสาร Excel หรือฐานข้อมูล SQL ตัวอย่างที่พบบ่อยของข้อมูลที่มีโครงสร้างรวมถึงข้อมูล GPS บันทึกการขาย ข้อมูลทางการเงินจากตลาดหุ้น เป็นต้น

2.2 ข้อมูลที่ไม่มีโครงสร้าง (Unstructured Data) หมายถึงข้อมูลที่ไม่มีความมาตรฐานที่แน่นอน มักใช้เวลามากขึ้นในการประมวลผลและวิเคราะห์ข้อมูลประเภทนี้ เช่น บันทึกเสียงหรือวิดีโอของความคิดเห็นของลูกค้า หรือโพสต์ทวิตเตอร์ เป็นต้น

2.3 ข้อมูลแบบกึ่งโครงสร้าง (Semi-structured Data) ยังไม่มีคำจำกัดความชัดเจนสำหรับข้อมูลประเภทนี้ แต่โดยทั่วไปมักจะประกอบด้วยข้อมูลที่มีทั้งรูปแบบโครงสร้างและแบบไม่มีโครงสร้างปรากฏ เช่น JSON, XML, CSV เป็นต้น สำหรับการประมวลผลและวิเคราะห์ข้อมูลประเภทนี้ทำได้ง่ายกว่าข้อมูลที่ไม่มีโครงสร้างเนื่องจากเครื่องมือการวิเคราะห์ในปัจจุบันสามารถแปลง อ่าน และประมวลผลข้อมูลแบบกึ่งโครงสร้างได้อย่างมีประสิทธิภาพ

3. Velocity (ความเร็ว) หมายถึงความเร็วที่ข้อมูลถูกสร้างขึ้น ข้อมูลขนาดใหญ่โดยทั่วไปถูกสร้าง จับ และประมวลผลในเวลาจริงหรือใกล้เคียงกับเวลาจริง สิ่งนี้สำคัญเนื่องจากมันช่วยให้องค์กรสามารถระบุปัญหาอย่างรวดเร็ว ปรับแผนการทำงาน และตอบสนองต่อสภาพแวดล้อมธุรกิจที่เปลี่ยนแปลงอย่างรวดเร็ว

4. Veracity (ความถูกต้อง) หมายถึงคุณภาพของข้อมูลในทางความถูกต้อง ความเชื่อถือได้ และความสามารถในการใช้งาน หากข้อมูลที่เก็บรวบรวมไม่ถูกต้อง กล่าวคือ ข้อมูลที่มีคุณภาพต่ำ ไม่ว่าจะมามีปริมาณมากเพียงใด ก็กลายเป็นข้อมูลที่ไร้ประโยชน์และไม่น่าเชื่อถือ (ขยะเข้าไปขยะออกมา (GIGO)) แม้ว่าเราจะพัฒนาเครื่องมือวิเคราะห์ที่ซับซ้อนและแข็งแกร่งอย่างไรก็ตาม ผลลัพธ์ที่ได้จะเชื่อถือไม่ได้ (อาจก่อให้เกิดความเสียหายได้) สำหรับการตัดสินใจที่มีประสิทธิภาพ ดังนั้น จึงเป็นสิ่งจำเป็นที่จะต้อง

หลีกเลี่ยงความล่าช้าในการเก็บข้อมูลและดำเนินการทำความสะอาดข้อมูลทุกครั้งก่อนนำข้อมูลเข้าสู่โมเดลเพื่อการวิเคราะห์ทุกครั้ง

- 5. Value (มูลค่า)** มูลค่าของข้อมูลขนาดใหญ่คือความสามารถในการเปลี่ยนข้อมูลเหล่านี้ให้เป็นข้อมูลเชิงลึก (insights) ที่มีคุณค่า ซึ่งสามารถนำไปใช้ในการตัดสินใจเชิงธุรกิจ การพัฒนาสินค้า หรือการปรับปรุงการบริการให้มีประสิทธิภาพมากยิ่งขึ้น

2.5 ข้อมูลขนาดใหญ่ในชีพพลายเซน

ข้อมูลขนาดใหญ่ในชีพพลายเซนดังที่กล่าวไว้แล้วข้างต้นว่าเป็นข้อมูลที่มีปริมาณมาก (volume) มีความหลากหลาย (variety) มีการเปลี่ยนแปลงอย่างรวดเร็ว (velocity) และมักมีความไม่แน่นอนในเชิงคุณภาพ (veracity) การจัดการข้อมูลขนาดใหญ่ในชีพพลายเซนสามารถช่วยให้ธุรกิจสามารถทำงานได้อย่างมีประสิทธิภาพมากขึ้นโดยการเพิ่มความโปร่งใสและความแม่นยำในการตัดสินใจ ประเภทของข้อมูลขนาดใหญ่ในชีพพลายเซนสามารถแบ่งออกเป็นประเภทต่างๆตามลักษณะการใช้งานได้ดังนี้:

1. ข้อมูลธุรกรรม (Transactional Data) ข้อมูลที่เกี่ยวข้องกับกิจกรรมต่าง ๆ ในชีพพลายเซน เช่น การสั่งซื้อสินค้า การจัดส่ง การผลิต การจัดเก็บ และการคืนสินค้า ข้อมูลนี้มักจะมาจากระบบ ERP หรือ WMS

2. ข้อมูลการปฏิบัติงาน (Operational Data) ข้อมูลที่เกี่ยวข้องกับกระบวนการทางกายภาพในชีพพลายเซน เช่น ข้อมูลการผลิต การบำรุงรักษาเครื่องจักร การจัดการคลังสินค้า และการขนส่ง

3. ข้อมูลเซ็นเซอร์และ IoT (Sensor and IoT Data) ข้อมูลที่ได้จากเซ็นเซอร์และอุปกรณ์ IoT ที่ติดตั้งอยู่ในชีพพลายเซน เช่น เซ็นเซอร์ตรวจจับอุณหภูมิและความชื้นในคลังสินค้า เซ็นเซอร์ตรวจสอบสภาพเครื่องจักรในโรงงาน ข้อมูลการสั่นสะเทือนของเครื่องจักรที่ใช้ในการผลิตเพื่อทำนายการบำรุงรักษา ข้อมูลการติดตามยานพาหนะในระหว่างการขนส่ง

4. ข้อมูลจากลูกค้าและตลาด (Customer and Market Data) ข้อมูลที่เกี่ยวข้องกับพฤติกรรมของลูกค้า ความต้องการของตลาด และแนวโน้มการบริโภค ข้อมูลนี้อาจมาจากหลายแหล่ง เช่น โซเชียลมีเดีย เว็บไซต์ของบริษัท การสำรวจตลาด หรือความคิดเห็นของลูกค้า ได้แก่ ข้อมูลการ

ค้นหาสินค้าในเว็บไซต์ของบริษัท ข้อมูลการรีวิวลสินค้าในโซเชียลมีเดีย ข้อมูลการทำแบบสำรวจความพึงพอใจของลูกค้า ข้อมูลยอดขายตามฤดูกาล

5. ข้อมูลสภาพอากาศและเหตุการณ์ภายนอก (Weather and External Event Data)

ข้อมูลที่เกี่ยวข้องกับสภาพอากาศ เหตุการณ์ทางเศรษฐกิจ หรือปัจจัยภายนอกอื่น ๆ ที่สามารถส่งผลกระทบต่อชัฟฟลายเซน ข้อมูลนี้มักมาจากแหล่งข้อมูลภายนอกเช่น หน่วยงานรายงานสภาพอากาศ หรือข่าวสารต่าง ๆ เช่น ข้อมูลสภาพอากาศที่อาจส่งผลต่อการขนส่งสินค้า ข้อมูลการระบาดของโรคที่อาจส่งผลกระทบต่อผลผลิต ข้อมูลการเมืองหรือเศรษฐกิจที่อาจส่งผลกระทบต่อความเสี่ยงในชัฟฟลายเซน

6. ข้อมูลการเงิน (Financial Data) ข้อมูลที่เกี่ยวข้องกับธุรกรรมทางการเงินในชัฟฟลายเซน เช่น ต้นทุนการผลิต ต้นทุนการขนส่ง ราคาวัตถุดิบ และข้อมูลการชำระเงิน ข้อมูลนี้มักมาจากระบบบัญชีหรือ ระบบERP

ดังที่กล่าวไว้แล้วข้างต้น เทคโนโลยีในปัจจุบันอำนวยความสะดวกในการเก็บข้อมูลขนาดใหญ่ในชัฟฟลายเซน จึงทำให้มีการนำข้อมูลเหล่านี้ไปใช้งานในแอปพลิเคชันต่าง ๆ อย่างมากมาย เช่น (1) การคาดการณ์ความต้องการสินค้าของลูกค้า โดยการนำข้อมูลในอดีต เช่น ข้อมูลการตลาดและข้อมูลจากภายนอก (เช่น สภาพอากาศ หรือ แนวโน้มเศรษฐกิจ) มาใช้ทำนายความต้องการสินค้าของลูกค้าในอนาคต นอกจากนั้นเนื่องจากกระแสนิยมที่เพิ่มมากขึ้นในปัจจุบันของสื่อสังคมออนไลน์ ทำให้จากเดิมที่ปกติใช้แต่ข้อมูลยอดขายมาพยากรณ์ความต้องการสินค้าอย่างเดียว เปลี่ยนมาเป็นการใช้ข้อมูลโซเชียลมีเดียมาเป็นปัจจัยร่วมที่สำคัญในการพยากรณ์ เช่น กระแสนิยมสินค้าในช่วงปีใหม่ที่ถูกโพสต์ลงในเพจ Facebook ของลูกค้า เป็นต้น (2) การเพิ่มประสิทธิภาพการขนส่ง โดยใช้ข้อมูลโลจิสติกส์ รวมถึงข้อมูลสภาพอากาศ สภาพการจราจร และข้อมูลเซ็นเซอร์ระบุตำแหน่งจากยานพาหนะ เพื่อกำหนดเส้นทางการขนส่งที่เร็วที่สุดและมีต้นทุนต่ำที่สุด (3) การจัดการสินค้าคงคลัง โดยการใช้ข้อมูลการขาย ข้อมูลการผลิต และข้อมูลเซ็นเซอร์ในคลังสินค้าเพื่อคาดการณ์และสั่งสินค้าล่วงหน้าเมื่อสินค้ากำลังจะหมด ส่งผลให้ลดการขาดแคลนสินค้าและลดต้นทุนการจัดเก็บ (4) การบำรุงรักษาเชิงคาดการณ์ โดยการใช้ข้อมูลเซ็นเซอร์จากเครื่องจักร (ที่วัดการสั่นของมอเตอร์และอุณหภูมิ) รวมไปถึงข้อมูลการบำรุงรักษาในอดีตเพื่อทำนายการเสียหายของเครื่องจักรและวางแผนการบำรุงรักษาเชิงป้องกันล่วงหน้า ทำให้ลดเวลาหยุดงานที่ไม่คาดคิดเนื่องจากเครื่องจักรมีปัญหาได้ หรือ (5) การใช้งานข้อมูลชัฟฟลายเซนที่กำลังนิยมในสถานการณ์ที่ไม่แน่นอนเกิดจาก สภาวะสงคราม เศรษฐกิจ โรคระบาดที่เกิดขึ้นในปัจจุบัน คือ การประเมินความเสี่ยงในชัฟฟลายเซน โดยการ

ใช้ข้อมูลภายนอก เช่น ข้อมูลลักษณะการกระจายตัวของโรคอุบัติใหม่ ข้อมูลสภาพอากาศ ข้อมูลการเมือง และข้อมูลทางเศรษฐกิจ เพื่อวางแผนรับมือหรือปรับกลยุทธ์อย่างทันท่วงที เพื่อลดปัญหาการหยุดชะงักของซัพพลายเชน (supply chain disruption) อันเป็นต้นเหตุสำคัญของการขาดแคลนวัตถุดิบทำให้ไม่สามารถผลิตสินค้าสู่ตลาด ส่งผลให้ธุรกิจจำเป็นต้องเลิกกิจการ

โดยสรุปข้อมูลขนาดใหญ่ในซัพพลายเชนมีบทบาทสำคัญในการปรับปรุงกระบวนการและการตัดสินใจในทุกขั้นตอน ตั้งแต่การคาดการณ์ความต้องการสินค้าไปจนถึงการจัดการความเสี่ยง ข้อมูลเหล่านี้ช่วยให้ธุรกิจสามารถตอบสนองต่อความเปลี่ยนแปลงในตลาดได้อย่างรวดเร็ว และเพิ่มประสิทธิภาพในการดำเนินงานของซัพพลายเชนโดยรวม

สรุป

ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล หมายถึงระบบซัพพลายเชนที่ใช้ข้อมูลในการตัดสินใจและดำเนินการในทุกขั้นตอน ตั้งแต่การจัดหาวัตถุดิบ การผลิต การจัดเก็บ ไปจนถึงการกระจายสินค้า การใช้ข้อมูลที่ถูกต้องและทันสมัยช่วยเพิ่มประสิทธิภาพ ลดต้นทุน และเพิ่มความพึงพอใจของลูกค้าในกระบวนการทั้งหมด การเปลี่ยนแปลงสู่ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล เกิดขึ้นจากปัจจัยหลายประการที่ผลักดันให้ธุรกิจต้องปรับตัวเพื่อความอยู่รอดและเพื่อการแข่งขันในตลาดที่มีความซับซ้อนและเปลี่ยนแปลงอย่างรวดเร็ว โดยปัจจัยสำคัญที่ทำให้ต้องเปลี่ยนเป็นซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล ได้แก่ ความซับซ้อนของโครงข่ายซัพพลายเชนที่มากขึ้นเนื่องจากความจำเป็นที่จะต้องทำงานร่วมกันกับซัพพลายเออร์หลายรายในหลายประเทศ ความคาดหวังที่สูงขึ้นของลูกค้าและความต้องการที่เปลี่ยนแปลงอย่างรวดเร็ว การแข่งขันอย่างรุนแรงที่เพิ่มขึ้นในตลาด ความต้องการในการลดต้นทุนและเพิ่มประสิทธิภาพ กฎระเบียบมาตรฐานที่เข้มงวดมากขึ้น หรือ การเติบโตของเทคโนโลยีดิจิทัลอย่างก้าวกระโดดในปัจจุบัน ซึ่งปัจจัยหลังสุดนี้มีบทบาทอย่างมากที่เป็นต้นกำเนิดของข้อมูลที่มีปริมาณมากขึ้นอย่างทวีคูณ หรือเรียกอีกอย่างว่า ข้อมูลขนาดใหญ่โดยมีคุณลักษณะที่สำคัญ 5 ประการ หรือ 5Vs ได้แก่ volume variety velocity veracity และ value เทคโนโลยีที่เหนี่ยวนำให้เกิดการเปลี่ยนแปลงเหล่านี้ได้แก่ ระบบ ERP WMS TMS GPS เทคโนโลยี IoT และ เซ็นเซอร์ หรือ ระบบ barcode RFID เป็นต้น โดยในปัจจุบันได้มีการนำข้อมูลขนาดใหญ่นี้ไปประยุกต์ใช้ประโยชน์ในซัพพลายเชน ได้แก่ การพยากรณ์ความต้องการสินค้าตามฤดูกาล การลดเวลาการขนส่งสินค้า การคาดการณ์การสั่งซื้อสินค้าคงคลังเมื่อสินค้ากำลังจะหมด การบำรุงรักษาเชิงคาดการณ์เพื่อทำการซ่อมบำรุงล่วงหน้าก่อนเครื่องจักรจะมีปัญหา การลดความเสี่ยงในการขนส่งจากการใช้ข้อมูลสภาพอากาศ ข้อมูลการเมืองและเศรษฐกิจ ผลลัพธ์ที่ได้จากการประยุกต์ใช้ข้อมูลขนาดใหญ่นี้ช่วยให้ธุรกิจสามารถตอบสนองต่อการเปลี่ยนแปลงที่รวดเร็วในตลาดและเพิ่มประสิทธิภาพในการดำเนินการของซัพพลายเชนโดยรวม

แบบฝึกหัด

1. ทำการสำรวจชัพพลายเซนของบริษัทการศึกษา เพื่อตรวจสอบแหล่งข้อมูลอะไรที่มีและข้อมูลที่แตกต่างกันเหล่านั้นจะถูกใช้โดยการจัดการชัพพลายเซนของบริษัทการศึกษาอย่างไร
2. อะไรคือปัจจัยที่ส่งผลกระทบต่อการแบ่งปันข้อมูลในชัพพลายเซน ลองยกตัวอย่างที่เห็นได้ชัดเจนพร้อมทั้งอธิบายมา 1 ตัวอย่าง
3. ยกตัวอย่างสาเหตุที่ทำให้เกิด bullwhip effect (BE) ในชัพพลายเซน พร้อมทั้งอธิบายถึงวิธีหรือแนวทางที่จะลดหรือหลีกเลี่ยง BE นี้
4. หาบริษัทการศึกษา และสำรวจระบบ IT/IS ที่ถูกนำมาใช้ในชัพพลายเซนของบริษัทดังกล่าว อธิบายถึงแหล่งข้อมูลและการใช้งานของข้อมูลเหล่านี้เพื่อการตัดสินใจที่มีประสิทธิภาพและช่วยปรับปรุงผลการดำเนินการ
5. จงอธิบายถึงการนำเทคโนโลยี blockchain มาใช้ในบริษัทการศึกษา และปัญหา อุปสรรค ที่เกิดขึ้น
6. คุณลักษณะของข้อมูลขนาดใหญ่ประกอบด้วยอะไรบ้าง จงอธิบายพร้อมยกตัวอย่างกรณีศึกษา
7. จงยกตัวอย่างอธิบายการใช้งานข้อมูลขนาดใหญ่กับบริษัทที่นักศึกษาทำงานอยู่
8. ถ้าบริษัทแห่งหนึ่งเป็นบริษัทขนาดกลางและเล็ก (SME) ต้องการที่จะเปลี่ยนองค์กรเป็นองค์กรที่ขับเคลื่อนด้วยข้อมูล ควรที่จะมีแนวทางเริ่มต้นอย่างไร จงอธิบายเป็นขั้นตอนอย่างละเอียด

เอกสารอ้างอิง

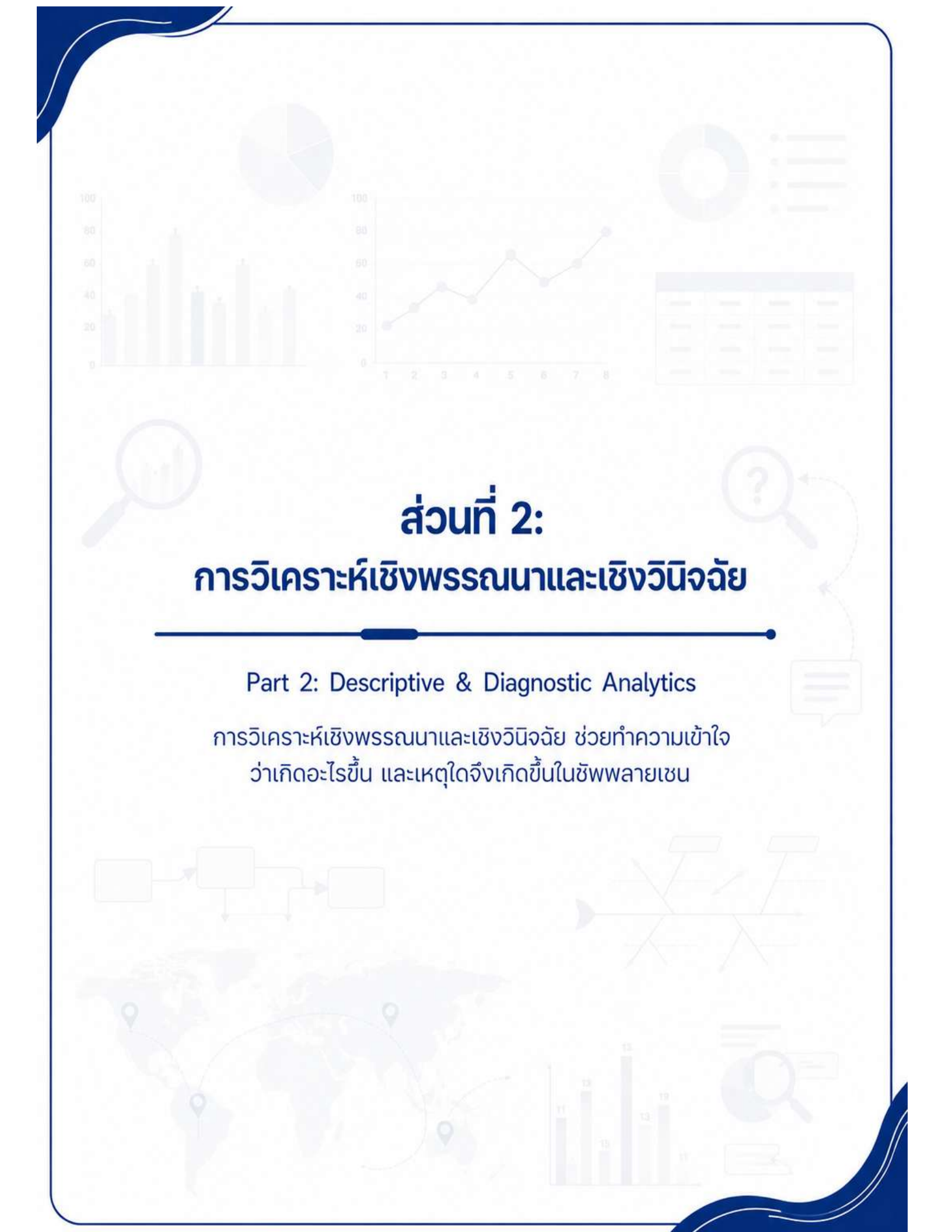
- Areerakulkan, N., & Pongpech, W. (2021). A Dempster-Shafer big data readiness assessment model. In *Proceedings of the 23rd International Conference on Enterprise Information Systems (ICEIS 2021)* (Vol. 2, pp. 581-585). SCITEPRESS. <https://doi.org/10.5220/0010506205810585>.
- Chanchaichujit, J., Balasubramanian, S., & Charmaine, N. S. M. (2020). A systematic literature review on the benefit-drivers of RFID implementation in supply chains and its impact on organizational competitive advantage. *Cogent Business & Management*, 7(1), 1818408. <https://doi.org/10.1080/23311975.2020.1818408>.
- Chen, X., Li, S., & Wang, X. (2020). Evaluating the effects of quality regulations on the pharmaceutical supply chain. *International Journal of Production Economics*, 230, Article 107770. <https://doi.org/10.1016/j.ijpe.2020.107770>.
- Dutta, P., Choi, T.-M., Somani, S., & Butala, R. (2020). Blockchain technology in supply chain operations: Applications, challenges and research opportunities. *Transportation Research Part E: Logistics and Transportation Review*, 142, 102067. <https://doi.org/10.1016/j.tre.2020.102067>.
- Feng, B., & Ye, Q. (2021). Operations management of smart logistics: A literature review and future research. *Frontiers of Engineering Management*, 8, 344-355. <https://doi.org/10.1007/s42524-021-0156-2>.
- Kache, F., & Seuring, S. (2017). Challenges and Opportunities of Digital Information at the Intersection of Big Data Analytics and Supply Chain Management. *International Journal of Operations & Production Management*, 37, 10-36. <https://doi.org/10.1108/ijopm-02-2015-0078>.
- Kamath, R. (2018). Food traceability on blockchain: Walmart's pork and mango pilots with IBM. *The Journal of the British Blockchain Association*, 1(1). [https://doi.org/10.31585/jbba-1-1-\(10\)2018](https://doi.org/10.31585/jbba-1-1-(10)2018).

เอกสารอ้างอิง (ต่อ)

- Koot, M., Mes, M. R. K., & Iacob, M. E. (2021). A systematic literature review of supply chain decision making supported by the Internet of Things and Big Data Analytics. *Computers & Industrial Engineering*, 154, 107076. <https://doi.org/10.1016/j.cie.2020.107076>.
- Lee, H. L., Padmanabhan, V., & Whang, S. (1997). Information distortion in a supply chain: The bullwhip effect. *Management Science*, 43(4), 546-558. <https://doi.org/10.1287/mnsc.43.4.546>.
- Liu, K. Y. (2022). *Supply Chain Analytics: Concepts, Techniques and Applications*. (1st ed.) Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-92224-5>.
- McFarlane, D., & Sheffi, Y. (2003). The impact of automatic identification on supply chain operations. *The International Journal of Logistics Management*, 14(1), 1-17. <https://doi.org/10.1108/09574090310806503>.
- Manthou, V., & Vlachopoulou, M. (2001). Bar-code technology for inventory and marketing management systems: A model for its development and implementation. *International Journal of Production Economics*, 71(1-3), 157-164. [https://doi.org/10.1016/S0925-5273\(00\)00115-8](https://doi.org/10.1016/S0925-5273(00)00115-8).
- Qureshi, M. R. N. M. (2022). Evaluating enterprise resource planning (ERP) implementation for sustainable supply chain management. *Sustainability*, 14(22), 14779. <https://doi.org/10.3390/su142214779>.
- Razak, G. M., Hendry, L. C., & Stevenson, M. (2023). Supply chain traceability: A review of the benefits and its relationship with supply chain resilience. *Production Planning & Control*, 34(11), 1114-1134. <https://doi.org/10.1080/09537287.2021.1983661>.
- Rodríguez-Espindola, O., Albores, P., & Brewster, C. (2018). Disaster preparedness in humanitarian logistics: A collaborative approach for resource management in floods. *European Journal of Operational Research*, 264(3), 978-993. <https://doi.org/10.1016/j.ejor.2017.01.021>.

เอกสารอ้างอิง (ต่อ)

- Singh, A., Gutub, A., Nayyar, A., & Khan, M. K. (2023). Redefining food safety traceability system through blockchain: Findings, challenges and open issues. *Multimedia Tools and Applications*, 82(14), 21243-21277. <https://doi.org/10.1007/s11042-022-14006-4>.
- Schoenherr, T., & Speier-Pero, C. (2015). *Data science, predictive analytics, and big data in supply chain management: Current state and future potential*. *Journal of Business Logistics*, 36(1), 120-132. <https://doi.org/10.1111/jbl.12082>.
- Tan, W. C., & Sidhu, M. S. (2022). Review of RFID and IoT integration in supply chain management. *Operations Research Perspectives*, 9, 100229. <https://doi.org/10.1016/j.orp.2022.100229>.
- Tarigan, Z. J. H., Siagian, H., & Jie, F. (2021). Impact of enhanced enterprise resource planning (ERP) on firm performance through green supply chain management. *Sustainability*, 13(8), 4358. <https://doi.org/10.3390/su13084358>.
- Unhelkar, B., Joshi, S., Sharma, M., Prakash, S., Mani, A. K., & Prasad, M. (2022). Enhancing supply chain performance using RFID technology and decision support systems in the industry 4.0—A systematic literature review. *International Journal of Information Management Data Insights*, 2(2), 100084. <https://doi.org/10.1016/j.jjime.2022.100084>



ส่วนที่ 2:

การวิเคราะห์เชิงพรรณนาและเชิงวินิจฉัย

Part 2: Descriptive & Diagnostic Analytics

การวิเคราะห์เชิงพรรณนาและเชิงวินิจฉัย ช่วยทำความเข้าใจ
ว่าเกิดอะไรขึ้น และเหตุใดจึงเกิดขึ้นในชีพพลายเซน

บทที่ 3

การวิเคราะห์ข้อมูลลูกค้า Customer Data Analytics

หัวข้อ

- 3.1 ลูกค้าในช่องทางขายและการทำความเข้าใจลูกค้า
- 3.2 การวิเคราะห์กลุ่มลูกค้า (Cohort analysis)
- 3.3 การวิเคราะห์ RFM
- 3.4 อัลกอริทึมการจัดกลุ่ม (Clustering Algorithms) เพื่อทำการแบ่งกลุ่มลูกค้า
- 3.5 การกำหนดนโยบายตามการแบ่งกลุ่มลูกค้า

การวิเคราะห์ข้อมูลลูกค้ามีความสำคัญอย่างมากสำหรับธุรกิจ เนื่องจากข้อมูลลูกค้าช่วยให้ผู้ประกอบการธุรกิจสามารถตัดสินใจและสร้างกลยุทธ์ที่มีประสิทธิภาพ เหตุผลที่การวิเคราะห์ข้อมูลลูกค้ามีความสำคัญได้แก่ (1) ทำให้เข้าใจพฤติกรรมและความต้องการของลูกค้า เพราะการวิเคราะห์ข้อมูลช่วยให้เข้าใจว่าลูกค้ามีพฤติกรรมอย่างไร ชอบอะไร และต้องการอะไร ทำให้สามารถพัฒนาผลิตภัณฑ์หรือบริการให้ตอบโจทย์ลูกค้าได้ดีขึ้น (2) ปรับปรุงการตลาด เนื่องจากข้อมูลลูกค้าสามารถช่วยสร้างแคมเปญการตลาดที่มีประสิทธิภาพมากขึ้น โดยสามารถกำหนดกลุ่มเป้าหมายได้อย่างแม่นยำ และเลือกวิธีการสื่อสารที่ตรงกับความต้องการของลูกค้าแต่ละแบบ (3) เพิ่มความพึงพอใจของลูกค้า การเข้าใจความต้องการและปัญหาของลูกค้าจากการวิเคราะห์ข้อมูลช่วยให้ผู้ประกอบการธุรกิจสามารถปรับปรุงบริการและผลิตภัณฑ์ได้ตามที่ลูกค้าคาดหวัง ส่งผลให้ลูกค้าพึงพอใจมากขึ้น (4) การรักษาลูกค้าเพราะข้อมูลที่วิเคราะห์ช่วยให้ผู้ประกอบการธุรกิจสามารถคาดการณ์ได้ว่าลูกค้ากลุ่มใดมีโอกาสที่จะยกเลิกการใช้บริการ และสามารถดำเนินการเชิงรุกเพื่อรักษาลูกค้าเหล่านั้นได้ (5) เพิ่มยอดขายจากการวิเคราะห์ข้อมูลช่วยให้ผู้ประกอบการธุรกิจสามารถเสนอผลิตภัณฑ์หรือบริการที่ตรงกับความต้องการของลูกค้าในเวลาที่เหมาะสม ทำให้มีโอกาสในการเพิ่มยอดขายได้มากขึ้น (6) ปรับปรุงประสิทธิภาพการดำเนินงานเพราะข้อมูลลูกค้าสามารถช่วยให้สามารถระบุปัญหาในการดำเนินงานและหาวิธีการปรับปรุง เพื่อให้การทำงานมีประสิทธิภาพและลดต้นทุน (7) การตัดสินใจที่ดีขึ้น ข้อมูลที่ได้รับจากการวิเคราะห์ช่วยให้ผู้บริหารมีข้อมูลที่จำเป็นในการตัดสินใจ และทำให้การตัดสินใจนั้นมีความแม่นยำและมีเหตุผลมากขึ้น ด้วยเหตุผลดังที่กล่าวมาแล้วการวิเคราะห์ข้อมูลลูกค้าจึงเป็นเครื่องมือสำคัญที่

ช่วยให้ธุรกิจสามารถพัฒนาตนเองและสร้างรายได้เปรียบในการแข่งขันได้อย่างยั่งยืน (Lemon & Verhoef, 2016; Wedel & Kannan, 2016; Fader, 2020)

3.1 ลูกค้าในซัพพลายเชนและการทำความเข้าใจลูกค้า

ในซัพพลายเชนลูกค้าคือใครบ้าง ? เป็นคำถามที่ไม่ได้มีคำตอบเดียวเสมอไป ลองจินตนาการว่าเวลาที่เราไปช้อปปิ้งในซูเปอร์มาร์เก็ต เราถือว่าเป็นลูกค้า หากเราซื้อธัญจากตัวแทนจำหน่ายรถเราก็ถือว่าเป็นลูกค้า เพราะฉะนั้นใครก็สามารถเป็นลูกค้าได้เมื่อมีการซื้อขายเกิดขึ้น เช่น ผู้ผลิตเป็นลูกค้าของ supplier tier 1 ในขณะที่ supplier tier 1 เป็นลูกค้า ของ supplier tier 2 ตามลำดับคำถามที่ตามมาคือแล้วลูกค้าที่แท้จริงหรือลูกค้าที่สำคัญที่สุดของซัพพลายเชนคือใคร ? คำตอบก็คือหากเรามองภาพรวมของซัพพลายเชน ลูกค้าสุดท้าย (end consumers) ถือว่าสำคัญที่สุด เพราะเป็นผู้ซื้อหรือผู้บริโภคสินค้าสำเร็จรูปของซัพพลายเชน ดังนั้นหากไม่มีผู้ซื้อสินค้าก็จะมีกระแสเงินสดไหลเข้าซัพพลายเชน เมื่อไม่มีกระแสเงินสดไหลเข้าธุรกิจต่าง ๆ ในซัพพลายเชนย่อมไม่สามารถอยู่รอดอย่างยั่งยืนได้ แต่โดยปกติแล้วธุรกิจต่าง ๆ ในซัพพลายเชนบางครั้งไม่ได้ตระหนักถึงความจริงข้อนี้โดยเฉพาะธุรกิจที่อยู่ทางต้นน้ำมักจะมองข้ามความสำคัญข้อนี้ไป ดังนั้นการที่ซัพพลายเชนจะประสบความสำเร็จนั้นสมาชิกทั้งหมดของซัพพลายเชนควรที่จะมีการวางกลยุทธ์ไปในทิศทางเดียวกันเพื่อสร้างความพึงพอใจให้กับลูกค้าสุดท้าย (Lambert & Cooper, 2000)

ลูกค้ามีบทบาทที่สำคัญในการจัดการซัพพลายเชน เนื่องจากวัตถุประสงค์หลักของการจัดการซัพพลายเชนคือการตอบสนองตามหรือเหนือความคาดหวังของลูกค้า (Mentzer et al., 2001) ถ้าไม่มีลูกค้าซัพพลายเชนก็ไม่เกิดขึ้นดังนั้นการจัดการซัพพลายเชนก็จะต้องมีความหมาย ธุรกิจประเภทใด ๆ ก็ตามควรพยายามให้บริการลูกค้าอย่างยอดเยี่ยมเท่าที่จะเป็นไปได้ เนื่องจากธุรกิจไม่สามารถอยู่รอดในระยะยาวหากไม่ให้บริการที่ดีกับลูกค้า นอกจากนั้นองค์กรที่ประสบความสำเร็จมักเข้าใจลูกค้าของตนได้ดีกว่าคู่แข่ง พวกเขาจะออกแบบซัพพลายเชนที่ตอบสนองอย่างรวดเร็ว มีประสิทธิภาพ และสร้างคุณค่าให้กับลูกค้าของพวกเขาให้มากที่สุด พวกเขาจะค้นหาลูกค้าที่มีนวัตกรรมและมีคุณภาพเพื่อให้บริการลูกค้าของพวกเขาด้วยสินค้าและบริการที่ดีที่สุด รวมไปถึงการออกแบบซัพพลายเชนที่ทนทานต่อความเสี่ยงและความขัดแย้งเพื่อให้บริการอย่างสม่ำเสมอและเชื่อถือได้ให้แก่ลูกค้าของพวกเขา จึงเป็นเรื่องที่ชัดเจนว่าการเข้าใจลูกค้าที่ดีกว่า จะนำมาซึ่งข้อได้เปรียบหลายประการสำหรับบริษัทหรือองค์กรนั้น ๆ ดังนั้นบริษัทส่วนใหญ่ที่ประสบความสำเร็จมักลงทุนเป็นเงินจำนวนมากในการทำวิจัยตลาดเพื่อให้เข้าใจลูกค้าของตนทั้งในด้านผลิตภัณฑ์ที่พวกเขาควรเสนอ คุณค่าอะไรที่พวกเขาควรมี และในกลุ่มตลาดไหนที่พวกเขาควรเน้น

ในตลาดที่แข่งขันอย่างรุนแรงในปัจจุบัน ธุรกิจทุกแห่งควรเข้าใจและสร้างความพึงพอใจลูกค้าของตนอย่างแท้จริงเพื่อแข่งขันกับคู่แข่ง ประสบการณ์ของลูกค้าที่มีความสำคัญเพิ่มขึ้นทำให้ต้องการซัพพลายเชนที่มุ่งเน้นลูกค้าเป็นศูนย์กลาง (customer-centric supply chain) คือซัพพลายเชนที่บริษัทเน้นสิ่งที่ลูกค้าต้องการจริง ๆ แล้วจึงจัดสรรเชิงกลยุทธ์ ทั้งทรัพยากร การประสานงาน และการพัฒนาความสามารถเพื่อที่จะบรรลุความคาดหวังของลูกค้า ไม่ว่าจะเป็นในอุตสาหกรรมใด ขนาดใด ภูมิภาคใด ใหม่หรือเก่า บริษัททุกแห่งควรให้ความสำคัญกับลูกค้า โดยเก็บความคาดหวังของลูกค้าไว้ในใจกลางของธุรกิจ และการดำเนินงานประจำวันของทุก ๆ วัน (Baldi et al., 2024; Rahman, 2024)

3.1.1 ซัพพลายเชนที่มุ่งเน้นลูกค้าเป็นศูนย์กลาง (Customer-centric Supply Chain)

การพัฒนาซัพพลายเชนที่เน้นลูกค้าเป็นศูนย์กลางอาจไม่่ง่ายอย่างที่คิด บริษัทจำเป็นต้องเปลี่ยนแปลงไม่เพียงแค่รูปแบบการติดต่อกับลูกค้าแบบเดิม ๆ เท่านั้น แต่ยังต้องปรับโครงสร้างซัพพลายเชนให้สอดคล้องกับกลยุทธ์เพื่อส่งมอบสิ่งที่ลูกค้าต้องการอย่างแท้จริง (Christopher, 2016; Chopra, 2019) นอกจากนี้สมาชิกในซัพพลายเชนควรมีส่วนร่วมทั้งทางตรงและทางอ้อม ตั้งแต่ผู้จัดหาวัตถุดิบ (ต้นน้ำ) ไปจนถึงลูกค้า(ปลายน้ำ) เพื่อสร้างประสบการณ์ที่ดีและยอดเยี่ยมให้กับลูกค้า วิธีการออกแบบและพัฒนาซัพพลายเชนที่เน้นลูกค้าเป็นศูนย์กลางนั้น สามารถสรุปเป็น 5 ขั้นตอนทั่วไปดังนี้

1. **ขั้นแรกการกำหนดลูกค้าที่แท้จริง** บริษัทควรมีการกำหนดลูกค้าที่แท้จริงในซัพพลายเชนของตนอย่างชัดเจน ดังที่กล่าวมาแล้ว สมาชิกในซัพพลายเชนควรให้ความสำคัญกับลูกค้าสุดท้ายที่ซื้อสินค้าสำเร็จรูป (Christopher, 2016) แทนที่จะมุ่งเน้นที่ส่วนใดส่วนหนึ่งของซัพพลายเชนหรือกระบวนการภายในของตนเอง ซึ่งเป็นสิ่งสำคัญอย่างยิ่งในการพัฒนาซัพพลายเชนที่เน้นลูกค้าเป็นศูนย์กลาง หากคุณสูญเสียตลาด คุณจะสูญเสียธุรกิจ และท้ายที่สุด ซัพพลายเชนทั้งหมดจะสูญเสียไปด้วย ดังนั้น เมื่อมีการกำหนดลูกค้าอย่างชัดเจนและสื่อสารไปทั่วทั้งซัพพลายเชนแล้ว ทุกคนที่เกี่ยวข้องจะทุ่มเทตนเองเพื่อสร้างซัพพลายเชนที่เน้นลูกค้าเป็นศูนย์กลางอย่างแท้จริง

2. **ขั้นตอนที่สองเข้าใจความต้องการที่แท้จริงของลูกค้า** หลังจากรู้ว่าใครคือลูกค้า ขั้นตอนต่อไปคือการเรียนรู้ความต้องการและความปรารถนาที่แท้จริงของพวกเขา บริษัทควรมีวิสัยทัศน์ที่ชัดเจนว่าเซกเมนต์ตลาดใดที่ควรมุ่งเน้น และข้อเสนอคุณค่าใดที่ควรยึดมั่นสำหรับเซกเมนต์ตลาดนั้นๆ

ตัวอย่างเช่น หากคุณเป็นบริษัทเริ่มต้นที่พัฒนาอุปกรณ์ติดตามกิจกรรมอัจฉริยะ (smart watch) คุณอาจต้องถามตัวเองดังตัวอย่างต่อไปนี้ (1) เราควรมุ่งเน้นที่เซ็กเมนต์ตลาดใด? (เช่น คนหนุ่มสาวหรือผู้สูงอายุ) (2) เราต้องการมอบคุณค่าอะไรให้กับลูกค้า? (เช่น ผลิตภัณฑ์คุณภาพสูงและฟังก์ชันที่หลากหลาย หรือผลิตภัณฑ์ราคาถูกลงกว่าคู่แข่งแต่มีคุณภาพยอมรับได้?) (3) ลูกค้าต้องการอะไรจากผลิตภัณฑ์จริง? (เช่น ความแม่นยำในการติดตามหรือความสามารถในการติดตามสุขภาพเพิ่มเติม เช่น การวัดความดันโลหิต?)

ปัจจุบัน บริษัทต่าง ๆ มีทางเลือกมากขึ้น เช่น การใช้การวิเคราะห์ข้อมูลขนาดใหญ่และการเรียนรู้ของเครื่อง (Ngai et al., 2009; Wedel & Kannan, 2016) เพื่อเรียนรู้เกี่ยวกับลูกค้าและคาดการณ์แนวโน้มในอนาคต ตัวอย่างเช่น บันทึกการซื้อของลูกค้า ประวัติการขาย รูปแบบการซื้อปีของของลูกค้า ข้อมูลประชากร และแม้แต่สภาพอากาศในภูมิภาคต่าง ๆ ล้วนเป็นข้อมูลที่มีประโยชน์ซึ่งสามารถใช้ในการวิเคราะห์และเข้าใจความต้องการที่แท้จริงของลูกค้า

3. ขั้นตอนที่สามแปลความต้องการเป็นคุณลักษณะของผลิตภัณฑ์ เมื่อระบุความต้องการที่แท้จริงของลูกค้าได้แล้ว ขั้นตอนที่สำคัญสำหรับบริษัทคือการแปลงความต้องการเหล่านั้นเป็นคุณลักษณะของผลิตภัณฑ์ (หรือบริการ) (Chopra, 2019) อย่างไรก็ตามนี้มักเป็นงานที่ทำหายเพราะองค์กรอาจ “หลงทางในการแปล” (เช่น ไม่สามารถตีความความต้องการที่แท้จริงและระบุคุณลักษณะที่เหมาะสมในการออกแบบผลิตภัณฑ์) หรือจำเป็นต้องปรับตัวเนื่องจากข้อจำกัดด้านเทคโนโลยีและทรัพยากร (เช่น เงินทุนและต้นทุน) บริษัทที่มีความสามารถอาจนำเสนอคุณลักษณะใหม่ในข้อเสนอของตนได้สำเร็จ และเอาชนะคู่แข่งเพื่อครองตลาดได้ ในขณะที่บริษัทอื่น ๆ ที่ไม่สามารถนำเสนอคุณลักษณะใหม่ได้อาจจะต้องลดต้นทุนและย้ายไปยังเซ็กเมนต์ที่แตกต่างกัน คุณลักษณะของผลิตภัณฑ์ทั่วไปบางประการ ได้แก่ ฟังก์ชันการใช้งาน คือ ฟังก์ชันต่าง ๆ ที่ผลิตภัณฑ์สามารถทำได้หรือมีอยู่ เช่น รูปลักษณ์การออกแบบภายนอกของผลิตภัณฑ์ (เช่น สี รูปทรง และพื้นผิว) ว่าดูมีสไตล์และดึงดูดลูกค้าหรือไม่ คุณภาพของผลิตภัณฑ์ว่าถูกผลิตตามมาตรฐานคุณภาพสูงหรือทำงานตามที่คาดหวังได้หรือไม่ เช่น มีความทนทานหรือไม่ ใช้งานได้ดีหรือไม่ ผลิตภัณฑ์มีความน่าเชื่อถือในการใช้งานและสามารถพึ่งพาได้เสมอหรือไม่ ผลิตภัณฑ์มีพร้อมจำหน่ายหรือไม่ การบริการต่อลูกค้าเป็นมิตรพอหรือไม่ ผลิตภัณฑ์ถูกส่งมอบถึงมือลูกค้าในความเร็วที่ยอมรับได้หรือไม่ หรือ ผลิตภัณฑ์มีความหลากหลายและสามารถปรับแต่งได้ตามความต้องการของลูกค้าหรือไม่

4. ขั้นตอนที่สี่ออกแบบกระบวนการซัพพลายเชน หลังจากแปลความต้องการเป็นคุณลักษณะของผลิตภัณฑ์แล้ว ขั้นตอนต่อไปและที่ทำหายที่สุดคือการออกแบบกระบวนการซัพพลายเชนที่

เหมาะสมเพื่อผลิตและส่งมอบผลิตภัณฑ์ที่ถูกต้องให้กับลูกค้าที่ถูกต้อง ตัวอย่างเช่น หากลูกค้าของคุณให้ความสำคัญกับคุณภาพมากที่สุด กระบวนการของซัพพลายเชนสามารถรับประกันการผลิตคุณภาพสูงได้หรือไม่ ตั้งแต่การจัดหาวัตถุดิบ การตรวจสอบประสิทธิภาพของผู้จัดหา ไปจนถึงการรับประกันคุณภาพภายในและการกระจายสินค้าสุดท้ายให้กับลูกค้า หากลูกค้าต้องการตัวเลือกและการปรับแต่งผลิตภัณฑ์มากขึ้น ซัพพลายเชนของคุณมีความคล่องตัวและยืดหยุ่นพอที่จะรับมือกับความแปรปรวนของความต้องการและตอบสนองอย่างรวดเร็วต่อคำขอต่าง ๆ ของลูกค้าหรือไม่ การออกแบบกระบวนการซัพพลายเชนที่เหมาะสมจำเป็นต้องมีการคัดเลือกและประเมินซัพพลายเออร์ที่ดี ตัวอย่างเช่น ซัพพลายเออร์ที่มีความสามารถช่วยทำให้มั่นใจถึงคุณภาพและเป้าหมายด้านประสิทธิภาพ การสร้างความสัมพันธ์แบบพันธมิตรระยะยาวและเชิงกลยุทธ์นำมาซึ่งประโยชน์หลายประการ เช่น การเพิ่มความไว้วางใจ การลดต้นทุน การแก้ปัญหาาร่วมกัน และการจัดหาสินค้าที่สม่ำเสมอเชื่อถือได้ ความสัมพันธ์ที่ดีกับซัพพลายเออร์ในซัพพลายเชนยังช่วยให้บริษัทที่เป็นศูนย์กลางสามารถเข้าถึงทรัพยากรและความสามารถภายนอกที่บริษัทยังไม่มี และอาจสร้างอุปสรรคให้กับคู่แข่งหรือผู้เข้ามาใหม่ในการเลียนแบบความสามารถที่สำคัญของบริษัทนั้นๆ

5. ขั้นตอนที่ทำให้สร้างระบบโลจิสติกส์ที่มีประสิทธิภาพ ถึงแม้เป็นขั้นตอนสุดท้ายแต่ไม่ได้สำคัญน้อยไปกว่าขั้นตอนอื่น ๆ แม้ว่าโลจิสติกส์จะเป็นส่วนสำคัญส่วนหนึ่งของกระบวนการของซัพพลายเชน แต่การพิจารณาจ้างผู้ให้บริการต่างหากอาจเป็นประโยชน์หากไม่ใช่ธุรกิจหลักขององค์กร สอดคล้องกับความจริงที่ว่าปัจจุบันบริษัทส่วนใหญ่จ้างบริการกระจายสินค้าและโลจิสติกส์กับผู้ให้บริการโลจิสติกส์บุคคลที่สาม (3PL) ในขั้นตอนนี้การพัฒนาเครือข่ายการกระจายสินค้าที่มีประสิทธิภาพจะสามารถช่วยปรับปรุงประสิทธิภาพการจัดส่ง ลดเวลารอคอยของลูกค้า และทำให้ลูกค้าพึงพอใจในที่สุด เช่น การใช้การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่อง บริษัทสามารถปรับเส้นทางการจัดส่งให้เหมาะสม บริหารจัดการกองยานพาหนะได้อย่างมีประสิทธิภาพ และติดตามตำแหน่งของผลิตภัณฑ์แบบเรียลไทม์ ด้วยระบบ AI และ IoT บริษัทสามารถตรวจสอบวิธีการจัดการเก็บรักษา ขนส่ง และส่งมอบผลิตภัณฑ์ถึงหน้าประตูบ้านของลูกค้าได้ตามระยะเวลาและคุณภาพที่ลูกค้าต้องการ (Esper et al., 2020)

3.2 การวิเคราะห์กลุ่มลูกค้า (Cohort Analysis)

การวิเคราะห์ cohort ถูกใช้ในวงการวิทยาศาสตร์สังคมมาเป็นเวลานาน เพื่อศึกษาประชากรศาสตร์ จิตวิทยา ระบาดวิทยา วิทยาการการเมือง และสังคมวิทยา เป็นต้น โดยทั่วไปแล้ว การวิเคราะห์กลุ่มลูกค้าเป็นเทคนิควิจัยเชิงปริมาณเพื่อศึกษากลุ่มคนที่มีลักษณะที่คล้ายคลึงกันและผลกระทบที่เกิดขึ้นต่อตัวแปรผลลัพธ์บางอย่าง ในช่วงปีหลังๆ เทคนิคนี้ได้รับความนิยมมากขึ้นในหัวข้อที่เกี่ยวข้องกับธุรกิจ เช่น การตลาดเพื่อเข้าใจพฤติกรรมของลูกค้าโดยการจัดหมวดหมู่ลูกค้าเป็น “cohort”

ในทางวิทยาศาสตร์สังคม cohort หมายถึงบุคคลหรือกลุ่มบุคคลที่มีลักษณะที่คล้ายกัน (เช่น เพศ อายุ หรือวัฒนธรรม) หรือมีประสบการณ์ที่เหมือนกันในช่วงเวลาที่ระบุ (เช่น กลุ่มที่สำเร็จการศึกษาในปีเดียวกัน) สำหรับธุรกิจ cohort มักจะหมายถึงกลุ่มลูกค้าที่จัดกลุ่มรวมกันตามช่วงเวลาที่เหมาะสมเจาะจง เช่น วัน เดือน หรือสัปดาห์ เมื่อพวกเขาทำการซื้อครั้งแรกจากบริษัท การจัดกลุ่มลูกค้าเป็น cohort จะช่วยให้ธุรกิจเข้าใจพฤติกรรมและลักษณะของกลุ่มลูกค้าแต่ละกลุ่มได้ดีขึ้น (Fedushko & Ustyianovych, 2022)

3.2.1 ขั้นตอนการวิเคราะห์กลุ่มลูกค้า

เพื่อทำการวิเคราะห์กลุ่มลูกค้าอย่างมีประสิทธิภาพ มีขั้นตอนบางอย่างที่ต้องปฏิบัติตาม โดยทั่วไป ซึ่งควรเริ่มต้นด้วยการกำหนดประเภทกลุ่มลูกค้าที่เหมาะสม กำหนดขนาดของกลุ่มลูกค้า จากนั้นเลือกตัวชี้วัดที่เหมาะสมเพื่อเปรียบเทียบระหว่างกลุ่มลูกค้า ขั้นตอนถัดไปคือการตั้งค่าช่วงวันที่สำหรับการวิเคราะห์กลุ่มลูกค้า เมื่อทุกอย่างพร้อมแล้ว เราสามารถดำเนินการวิเคราะห์กลุ่มลูกค้าจริงได้ สุดท้ายเราสามารถแสดงผลและตีความผลลัพธ์ตามความเหมาะสม สรุปขั้นตอนการวิเคราะห์กลุ่มลูกค้าแสดงดังรายละเอียดดังต่อไปนี้

1. **กำหนดประเภทกลุ่มลูกค้าที่เหมาะสม** ขึ้นอยู่กับคำถามเฉพาะที่ต้องตอบ ขั้นตอนแรกคือการตัดสินใจว่าคุณต้องการสำรวจกลุ่มลูกค้าประเภทใดในการวิเคราะห์ ตัวอย่างเช่น การจัดกลุ่มตามอายุ เพศ รายได้ หรือช่วงเวลาของลูกค้าทำการซื้อครั้งแรก โดยประเภทหลังสุดจะเป็นประเภทที่ใช้กันมากที่สุดสำหรับการวิเคราะห์กลุ่มลูกค้าปัจจุบัน ยกตัวอย่างของร้านทำผม เราอาจจำแนกกลุ่มลูกค้าตามเวลาที่ลูกค้าเข้ามาใช้บริการเป็นครั้งแรก
2. **กำหนดขนาดของกลุ่มลูกค้า** เมื่อกำหนดประเภทกลุ่มลูกค้าแล้ว ขั้นตอนต่อไปคือการกำหนดขนาดของกลุ่มลูกค้า ขนาดในที่นี้หมายถึงช่วงหรือระยะของข้อมูลที่เราใช้ในการจัด

กลุ่มลูกค้า ตัวอย่างเช่น ช่วงอายุของแต่ละกลุ่มลูกค้า (เช่น 20-29 หรือ 30-39 เป็นต้น) และที่พบบ่อยที่สุดคือลูกค้าที่ได้รับการจัดกลุ่มตามสัปดาห์ เดือน หรือไตรมาส เป็นต้น ยกตัวอย่างเช่น การจัดกลุ่มลูกค้าตามเดือนอาจมีความเหมาะสมหากเราทราบว่าโดยเฉลี่ยแล้วลูกค้าจะมาใช้บริการร้านทำผมทุกเดือน

3. **เลือกตัวชี้วัด** ขั้นตอนในการเลือกค่าที่เฉพาะเจาะจงที่ต้องการเปรียบเทียบระหว่างลูกค้ากลุ่มต่าง ๆ เช่น อัตราการรักษาลูกค้าและจำนวนเงินที่ใช้จ่าย โดยทั่วไปการวิเคราะห์กลุ่มลูกค้าจะตรวจสอบอัตราการคงอยู่ของลูกค้าเป็นตัวชี้วัดหลัก (Ascarza et al., 2018) ในตัวอย่างของร้านทำผม เราสามารถตรวจสอบอัตราการคงอยู่ของลูกค้ารายเดือนในช่วงเวลาที่กำหนดได้
4. **ตั้งค่าช่วงวันที่** ที่ต้องการตรวจสอบด้วยการวิเคราะห์กลุ่มลูกค้า ตัวอย่างเช่น เราสามารถเปรียบเทียบกลุ่มลูกค้าสำหรับ 12 สัปดาห์ที่ผ่านมา ยกตัวอย่างเช่น สมมติว่าบริษัทได้ดำเนินธุรกิจมาเป็นเวลาหนึ่งปี เราอาจต้องการตรวจสอบอัตราการคงอยู่ของลูกค้ารายเดือนเทียบกับกลุ่มลูกค้าในช่วง 12 เดือน
5. **ดำเนินการวิเคราะห์กลุ่มลูกค้า** เมื่อทุกอย่างพร้อมแล้ว เราสามารถดำเนินการวิเคราะห์กลุ่มลูกค้าได้
6. **การแสดงผลและตีความผลลัพธ์** เป็นเรื่องปกติที่จะใช้การแสดงผลข้อมูลเพื่อแสดงผลการวิเคราะห์กลุ่มลูกค้า อย่าลืมมุ่งเน้นที่คำถามเริ่มต้นที่ต้องการตอบและตรวจสอบให้แน่ใจว่าการตีความมีความสมเหตุสมผล (Pereira et al., 2025)

ตัวอย่าง 3.1 ตัวอย่างการวิเคราะห์กลุ่มลูกค้าใน Microsoft Excel

ทำการเก็บข้อมูลจำนวนลูกค้าที่ Subscribe บริการ Sass Business ราคา \$50 ต่อเดือน เป็นเวลา 1 ปี ข้อมูลแสดงดังภาพที่ ต่อไปนี้

Month	Start Month	0	1	2	3	4	5	6	7	8	9	10	11
0	Jan	35											
1	Feb	35	18										
2	Mar	30	18	22									
3	Apr	26	16	22	21								
4	May	25	15	21	21	26							
5	Jun	24	13	19	20	26	20						
6	Jul	22	13	17	18	25	20	22					
7	Aug	22	12	17	17	23	19	22	22				
8	Sep	20	11	15	16	19	19	20	21	19			
9	Oct	18	10	11	14	18	19	20	19	19	23		
10	Nov	15	9	10	9	17	16	19	17	19	23	12	
11	Dec	9	9	9	9	15	14	18	13	17	23	12	16

ภาพที่ 3.1 จำนวนลูกค้าที่ Subscribe บริการ Sass Business ราคา \$50 ต่อเดือน

ที่มา: ภาพโดยผู้แต่ง

หมายเหตุ: วิธีการอ่านข้อมูลลูกค้าสมัครใหม่ครั้งแรกเดือน ม.ค. 35 ราย ถึงเดือน ก.พ. เหลือเท่าเดิมที่ 35 ราย ถึงเดือน มี.ค. เหลือ 30 ราย ไปจนถึงเดือน ธ.ค. เหลือ 9 ราย

2. ทำการแปลงจำนวนลูกค้าเป็น รายรับสุทธิ โดยใช้สมการที่ 3.1

$$\text{รายรับสุทธิ} = \text{จำนวนลูกค้ารายเดือน} \times \text{ค่าบริการรายเดือน} \quad (3.1)$$

ตย. รายรับสุทธิของ (Jan) = $35 \times 50 = \$1,750$ ต่อเดือน

หารายรับสุทธิของทั้งตารางจะได้ผลลัพธ์แสดงดังภาพที่ 3.2 ต่อไปนี้

Month	Start Month	0	1	2	3	4	5	6	7	8	9	10	11
0	Jan	1750											
1	Feb	1750	900										
2	Mar	1600	900	1100									
3	Apr	1500	800	1250	1050								
4	May	1350	700	1150	1050	1300							
5	Jun	1200	700	1100	850	1250	1000						
6	Jul	1100	700	900	800	1200	1000	1100					
7	Aug	1100	700	850	750	1150	1000	1100	1100				
8	Sep	1000	600	750	650	1000	1000	1100	1100	950			
9	Oct	800	400	600	550	900	900	1050	1000	950	1150		
10	Nov	600	350	650	500	750	800	950	800	950	1150	600	
11	Dec	450	250	400	350	700	750	850	750	850	1100	600	800

ภาพที่ 3.2 รายรับสุทธิต่อเดือน

ที่มา: ภาพโดยผู้แต่ง

3. คำนวณรายรับสะสมตลอดช่วงชีวิตของกลุ่มลูกค้าได้ผลลัพธ์แสดงดังภาพที่ 3.3 ต่อไปนี้

Month	Start Month	0	1	2	3	4	5	6	7	8	9	10	11
0	Jan	1750											
1	Feb	3500	900										
2	Mar	5100	1800	1100									
3	Apr	6600	2600	2350	1050								
4	May	7950	3300	3500	2100	1300							
5	Jun	9150	4000	4600	2950	2550	1000						
6	Jul	10250	4700	5500	3750	3750	2000	1100					
7	Aug	11350	5400	6350	4500	4900	3000	2200	1100				
8	Sep	12350	6000	7100	5150	5900	4000	3300	2200	950			
9	Oct	13150	6400	7700	5700	6800	4900	4350	3200	1900	1150		
10	Nov	13750	6750	8350	6200	7550	5700	5300	4000	2850	2300	600	
11	Dec	14200	7000	8750	6550	8250	6450	6150	4750	3700	3400	1200	800

ภาพที่ 3.3 รายรับสะสมตลอดช่วงชีวิต

ที่มา: ภาพโดยผู้แต่ง

4. คำนวณรายรับตลอดช่วงชีวิตของกลุ่มลูกค้า โดยใช้สมการที่ 3.2

$$\text{รายรับตลอดช่วงชีวิตของลูกค้า} = \text{รายรับสะสม} / \text{จำนวนลูกค้าแรกเข้า} \quad (3.2)$$

ตย. รายรับตลอดช่วงชีวิตของลูกค้า (Jan) = $1750 / 35 = \$50$ ต่อเดือน

ทำทั้งตารางจะได้ผลลัพธ์แสดงดังภาพที่ 3.4 ต่อไปนี้

Month	Start Month	0	1	2	3	4	5	6	7	8	9	10	11
0	Jan	50											
1	Feb	100	50										
2	Mar	146	100	50									
3	Apr	189	144	107	50								
4	May	227	183	159	100	50							
5	Jun	261	222	209	140	98	50						
6	Jul	293	261	250	179	144	100	50					
7	Aug	324	300	289	214	188	150	100	50				
8	Sep	353	333	323	245	227	200	150	100	50			
9	Oct	376	356	350	271	262	245	198	145	100	50		
10	Nov	393	375	380	295	290	285	241	182	150	100	50	
11	Dec	406	389	398	312	317	323	280	216	195	148	100	50

ภาพที่ 3.4 รายรับตลอดช่วงชีวิต

ที่มา: ภาพโดยผู้แต่ง

5. คำนวณมูลค่าตลอดช่วงชีวิตของกลุ่มลูกค้า (customer lifetime value, CLV) โดยใช้สมการที่ 3.3

$$CLV = \text{รายรับตลอดช่วงชีวิตของลูกค้า} \times \% \text{ กำไรขั้นต้น (Gross Margin)} \quad (3.3)$$

$$CLV (\text{Jan}) = 50 \times 65\% = \$33 \text{ ต่อเดือน ทำทั้งตารางจะได้ผลลัพธ์}$$

แสดงดังภาพที่ 3.5 ต่อไปนี้

Month	Start Month	0	1	2	3	4	5	6	7	8	9	10	11
0	Jan	33											
1	Feb	65	33										
2	Mar	95	65	33									
3	Apr	123	94	69	33								
4	May	148	119	103	65	33							
5	Jun	170	144	136	91	64	33						
6	Jul	190	170	163	116	94	65	33					
7	Aug	211	195	188	139	123	98	65	33				
8	Sep	229	217	210	159	148	130	98	65	33			
9	Oct	244	231	228	176	170	159	129	95	65	33		
10	Nov	255	244	247	192	189	185	157	118	98	65	33	
11	Dec	264	253	259	203	206	210	182	140	127	96	65	33

ภาพที่ 3.5 มูลค่าของลูกค้าตลอดช่วงชีวิต

ที่มา: ภาพโดยผู้แต่ง

CAC คือ ต้นทุนการได้มาของลูกค้า (customer acquisition cost - CAC): ค่าใช้จ่ายเฉลี่ยที่เกิดขึ้นในการดึงดูดลูกค้าใหม่ ซึ่งรวมถึงค่าใช้จ่ายทั้งหมดในการตลาดและการขายหารด้วยจำนวนลูกค้าใหม่ที่ได้มาในช่วงเวลาที่กำหนด

$$CAC = \frac{\text{ค่าใช้จ่ายทั้งหมดด้านการขายและการตลาด (Total Sales and Marketing Expenses)}}{\text{จำนวนลูกค้าใหม่ที่ได้มา (Number of New Customers Acquired)}} \quad (3.4)$$

สำหรับตัวอย่างนี้ถ้าสมมติให้กำหนดต้นทุนทางการตลาดทั้งหมด CAC เท่ากับ \$115 เมื่อทำการวิเคราะห์ค่า CLV พบว่าบริษัทมีผลการดำเนินงานอยู่ในระดับดีและสามารถอยู่รอดได้ เพราะ CLV ของลูกค้าส่วนใหญ่สูงกว่า CAC ภายในไม่กี่เดือน ตัวอย่างเช่น ลูกค้ากลุ่ม ม.ค. คืนทุนประมาณเดือนที่ 4 เพราะ CLV = \$123 มากกว่า CAC = \$115 เช่นเดียวกัน ลูกค้ากลุ่ม ก.พ. คืนทุนประมาณเดือนที่ 4 เพราะ CLV = \$119 มากกว่า CAC = \$115 เช่นเดียวกัน ลูกค้ากลุ่ม พ.ค. คืนทุนประมาณเดือนที่ 4 เพราะ CLV = \$123 มากกว่า CAC = \$115 เป็นต้น หรือกล่าวอีกนัยหนึ่งว่า ธุรกิจนี้น่าสนใจเพราะ CLV ของลูกค้าส่วนใหญ่สูงกว่า CAC ภายในไม่กี่เดือนแสดงว่าการหาลูกค้ามีความคุ้มค่า และหลังจากคืนทุนแล้วรายได้ที่เกิดขึ้นต่อไปจะกลายเป็นกำไรขั้นต้นของธุรกิจ

3.3 การวิเคราะห์ RFM

การวิเคราะห์ RFM เป็นเทคนิคการแบ่งกลุ่มลูกค้าที่มีประโยชน์ ซึ่งถูกนำมาใช้กันอย่างแพร่หลายในงานวิจัยทางการตลาด เทคนิคนี้สามารถช่วยในการจัดกลุ่มลูกค้าออกเป็นคลัสเตอร์ต่างๆ ทำให้ธุรกิจสามารถมุ่งเป้าหมายไปยังกลุ่มลูกค้าที่เฉพาะเจาะจงด้วยการสื่อสารที่ตรงจุด แคมเปญการมีส่วนร่วม และการดูแลพิเศษที่มีความเกี่ยวข้องมากขึ้นสำหรับแต่ละกลุ่ม (Wang et al., 2024; Kumar et al., 2025) ในส่วนนี้ เราจะแนะนำพื้นฐานของการวิเคราะห์ RFM และใช้ตัวอย่างในการสาธิตวิธีการทำการวิเคราะห์ด้วย Microsoft Excel

3.3.1 การวิเคราะห์ RFM

RFM ย่อมาจาก recency (ซื้อครั้งล่าสุดเมื่อไร) frequency (ความถี่ในการซื้อ) และ monetary (มูลค่าเงินที่ซื้อ) แต่ละตัวมีค่าสำคัญในการวิเคราะห์ลูกค้า (ดูภาพที่ 3.6)



ภาพที่ 3.6 ความหมายของ RF

ที่มา: ภาพโดยผู้แต่ง

ทั้งสามเมตริก RFM (recency | frequency | monetary) เป็นตัวชี้วัดที่สำคัญในการกำหนดพฤติกรรมของลูกค้า เนื่องจากสามารถสะท้อนถึงการคงอยู่ของลูกค้า การมีส่วนร่วม และมูลค่าตลอดชีวิตของลูกค้าได้

1. **Recency (ซื้อครั้งล่าสุดเมื่อไร)** หมายถึงว่าลูกค้าได้ทำการซื้อสินค้าครั้งล่าสุดเมื่อใด โดยวัดเวลาที่ผ่านไปนับตั้งแต่การสั่งซื้อหรือทำธุรกรรมครั้งสุดท้ายของลูกค้ากับธุรกิจ คะแนน recency สูงมักหมายความว่าลูกค้าได้พิจารณาธุรกิจของคุณในทางบวกในการตัดสินใจซื้อสินค้าแทนที่จะเลือกคู่แข่ง และลูกค้าอาจจะตอบสนองต่อการสื่อสารและแคมเปญการมีส่วนร่วมจากธุรกิจของคุณมากกว่า
2. **Frequency (ความถี่)** หมายถึงความบ่อยครั้งที่ลูกค้าทำการซื้อสินค้า สามารถวัดได้จากจำนวนคำสั่งซื้อทั้งหมดจากลูกค้าในช่วงเวลาหนึ่ง หรือช่วงเวลาเฉลี่ยระหว่างคำสั่งซื้อ คะแนนความถี่สูงมักหมายความว่าลูกค้ามีการมีส่วนร่วมและภักดีต่อธุรกิจของคุณ
3. **Monetary (การใช้จ่าย)** หมายถึงจำนวนเงินที่ลูกค้าใช้จ่ายทั้งหมดหรือเฉลี่ยในช่วงเวลาหนึ่ง คะแนนการใช้จ่ายสูงมักหมายความว่าลูกค้าเป็นผู้ใช้จ่ายมากที่ธุรกิจของคุณควรให้ความสนใจอย่างใกล้ชิด

การวิเคราะห์ RFM ผสมผสานคุณลักษณะลูกค้าทั้งสามเมตริกและจัดกลุ่มลูกค้าด้วยการจัดอันดับตัวเลขของสามเมตริกนี้ทำให้ได้มุมมองที่สมคูลมากขึ้น โดยทั่วไปลูกค้าที่มีคะแนนรวมสูงสุดจะถือเป็นลูกค้าที่ดีที่สุด ทั้งนี้ขึ้นอยู่กับบริบทของธุรกิจ การวิเคราะห์ RFM เป็นวิธีการที่นิยมใช้สำหรับการแบ่งกลุ่มลูกค้า เนื่องจากง่ายต่อการดำเนินการ และผลลัพธ์สามารถตีความและเข้าใจได้ง่าย โดยเฉพาะอย่างยิ่ง RFM ช่วยให้ตอบคำถามที่สำคัญได้แก่ (1) ใครคือลูกค้าที่มีคุณค่ามากที่สุด? (2) ใครมีศักยภาพที่จะกลายเป็นลูกค้าที่มีคุณค่ามากขึ้น? (3) ลูกค้าคนใดที่สามารถรักษาไว้ได้? (4) ลูกค้าคนใดมีแนวโน้มที่จะตอบสนองต่อแคมเปญการมีส่วนร่วม?

3.3.2 ขั้นตอนในการวิเคราะห์ RFM

ขั้นตอนในการวิเคราะห์ RFM มีรายละเอียดต่อไปนี้

1. กำหนดขอบเขตการวิเคราะห์ RFM นี่เป็นขั้นตอนในการกำหนดช่วงเวลาที่จะครอบคลุมในการวิเคราะห์ RFM ตัวอย่างเช่น เราสามารถทำการวิเคราะห์ RFM สำหรับช่วง 3 หรือ 6 เดือนที่ผ่านมา หรือแม้กระทั่งนานกว่านั้น แน่นอนว่าขึ้นอยู่กับความพร้อมของข้อมูล ขั้นตอนนี้มีความสำคัญเนื่องจากจะเป็นการตั้งจุดอ้างอิงสำหรับการคำนวณคะแนน $R | F | M$

2. การคำนวณค่า $R | F | M$ เมื่อกำหนดขอบเขตแล้ว คุณสามารถคำนวณค่า R F และ M ได้ตามลำดับ ดังนี้:

1) R (Recency) เพื่อคำนวณค่า R เพียงแค่หวันที่ซื้อสินค้าล่าสุดของลูกค้าออกจากวันที่ปัจจุบัน ความแตกต่างที่ได้จะเป็นค่าของ R (เช่น 1, 10 หรือ 120 วัน) ค่าที่น้อยกว่าหมายถึงลูกค้าที่เพิ่งทำการซื้อไม่นาน

2) F (Frequency) ในการคำนวณค่า F สามารถคำนวณจำนวนคำสั่งซื้อทั้งหมดจากลูกค้าในช่วงเวลาที่กำหนด หรือใช้เวลาระหว่างคำสั่งซื้อเฉลี่ย สำหรับตัวเลือกแรกยิ่งค่ามากยิ่งดี ส่วนตัวเลือกที่สอง ยิ่งค่าน้อยยิ่งดี

3) M (Monetary) สำหรับค่า M สามารถคำนวณจำนวนเงินทั้งหมดที่ลูกค้าใช้จ่ายในช่วงเวลานั้น หรือใช้ค่าเฉลี่ยของเงินที่ใช้จ่ายต่อคำสั่งซื้อ

การคำนวณเหล่านี้ช่วยให้สามารถวิเคราะห์และแบ่งกลุ่มลูกค้าได้ตามพฤติกรรมการซื้อขายของพวกเขา ตัวอย่างการคำนวณค่า RFM แสดงดังตารางต่อไปนี้

ตารางที่ 3.1 ตารางตัวอย่างการคำนวณค่า RFM

Customer ID	Recency (วัน)	Frequency (เวลา)	Monetary (ผลรวม)
101356	6	16	27,000
101365	8	30	123,000
101366	19	7	53,435
101367	94	2	1,200
101368	3	5	4,350

ที่มา: จัดทำโดยผู้แต่ง

3. การแปลงค่า R | F | M เป็นคะแนนตามสเกล ขั้นตอนนี้เป็นการแปลงค่า R | F | M สามค่าเป็นคะแนนตามสเกล ธุรกิจต่าง ๆ อาจใช้สเกลที่แตกต่างกันไป แต่สเกลที่พบบ่อยที่สุดคือ สเกล 1-5 ในการกำหนดสเกลนี้ เราสามารถตั้งค่าช่วงคงที่หรือคำนวณจากควอนไทล์ตามค่า R | F | M ทั้งสามค่า ตัวอย่างเช่น เราสามารถแบ่งการกระจายของค่า R ออกเป็นห้ากลุ่มเท่ากันโดยใช้ควอนไทล์สี่ค่า (เช่น [0.2, 0.4, 0.6, 0.8]) และค่าที่ต่ำกว่าค่าควอนไทล์แรก 0.2 จะถูกกำหนดให้คะแนน 5 (ข้อสังเกตค่า R ที่น้อยหมายถึงการซื้อหรือเยี่ยมชมล่าสุดจากลูกค้า หรืออีกนัยหนึ่งค่า R ก่อนการแปลงสเกลยิ่งน้อยยิ่งดี ดังนั้นคะแนน R ที่สูงสุดคือ 5 หมายถึงช่วงควอนไทล์แรก (≤ 0.2)) ในทำนองเดียวกัน เราสามารถคำนวณควอนไทล์สำหรับค่า F และ M และจากนั้นกำหนดคะแนนตามสเกลให้กับลูกค้าแต่ละคน

4. การใช้สูตร RFM สำหรับคะแนนรวม RFM เราสามารถหยุดที่ขั้นตอนก่อนหน้านี้ และใช้คะแนนตามสเกล R | F | M แยกกันสำหรับการแบ่งกลุ่มลูกค้า (เช่น 1-1-1 สำหรับกลุ่มต่ำ และ 5-5-5 สำหรับกลุ่มสูง) หรือธุรกิจสามารถสร้างสูตรของตนเองเพื่อหาคะแนนรวม RFM สำหรับลูกค้าแต่ละรายได้

ตัวอย่างเช่น เราสามารถใช้ค่าเฉลี่ยของคะแนนตามสเกล RFM ทั้งสาม หรือใช้ผลรวมก็ได้ ขึ้นอยู่กับลักษณะของธุรกิจ เราอาจกำหนดน้ำหนักที่แตกต่างกันสำหรับตัวแปร R | F | M เพื่อให้ได้คะแนน RFM ที่เหมาะสมกับธุรกิจของคุณมากขึ้น เช่น มูลค่าการใช้จ่ายต่อการทำธุรกรรมอาจมีความสำคัญมากกว่าสำหรับร้านขายเปียโนเมื่อเทียบกับร้านสะดวกซื้อที่ความใหม่และความถี่อาจมีความสำคัญมากกว่า

ตัวอย่าง 3.2 ตัวอย่างการคำนวณคะแนนรวม RFM

1. ใช้ค่าเฉลี่ยของคะแนน R, F, M:

$$\text{ค่าเฉลี่ยของคะแนน RFM score} = (R+F+M)/3 \quad (3.5)$$

2. ใช้ผลรวมของคะแนน R, F, M:

$$\text{ผลรวมของคะแนน RFM score} = R+F+M \quad (3.6)$$

3. ใช้สูตรที่มีน้ำหนักสำหรับตัวแปร R, F, M:

$$\text{คะแนนถ่วงน้ำหนักของ RFM score} = w_R \times R + w_F \times F + w_M \times M \quad (3.7)$$

โดยที่ w_R , w_F , และ w_M คือน้ำหนักที่กำหนดให้กับค่า R, F, M ตามความสำคัญของแต่ละตัวแปรในบริบทธุรกิจ

4. ตัวอย่างการกำหนดน้ำหนัก

- 1) สำหรับร้านขายเปียโน: (ตัวเงินสำคัญที่สุด)

$$w_R=0.2, w_F=0.3, w_M=0.5$$

- 2) สำหรับร้านสะดวกซื้อ: (เวลาซื้อล่าสุด และ ความถี่สำคัญที่สุด โดยมีค่าน้ำหนักเท่ากัน)

$$w_R=0.4, w_F=0.4, w_M=0.2$$

ตัวอย่างการแปลงค่า RFM เป็นค่าคะแนนแสดงดังตารางต่อไปนี้

ตารางที่ 3.2 ตารางตัวอย่างการแปลงค่า RFM

Customer ID	Recency (วัน)	Frequency (เวลา)	Monetary (ผลรวม)	RFM score (ค่าเฉลี่ย)
101356	5	3	4	4
101365	5	4	5	4.67
101366	4	1	4	3
101367	1	1	1	1
101368	5	1	1	2.33

ที่มา: จัดทำโดยผู้แต่ง

5. แบ่งกลุ่มลูกค้า ขั้นตอนสุดท้ายคือการสร้างการแบ่งกลุ่มลูกค้าตามคะแนน RFM ขึ้นอยู่กับลักษณะของธุรกิจ คุณสามารถตัดสินใจเลือกวิธีที่เหมาะสมที่สุดในการแบ่งลูกค้าออกเป็นกลุ่มต่างๆ ตัวอย่างการแบ่งกลุ่มลูกค้าแสดงดังตารางต่อไปนี้

ตารางที่ 3.3 ตารางตัวอย่างการแบ่งกลุ่มลูกค้าด้วยการใช้ค่า RFM

กลุ่ม	คะแนน RFM	ลักษณะเฉพาะกลุ่ม
Diamond	4.5-5.0	ลูกค้าที่ดีที่สุดคือลูกค้าที่ซื้อสินค้าหรือบริการล่าสุดเร็วที่สุด (recency) ซื้อบ่อยที่สุด (frequency) และเป็นผู้ใช้จ่ายเงินมากที่สุด (monetary)
Platinum	4.0-4.5	ลูกค้าที่มีมูลค่าที่สุดอันดับสอง ที่มีศักยภาพที่เลื่อนไปเป็นลูกค้ากลุ่มที่ดีที่สุด
Gold	3.5-4.0	ลูกค้ามูลค่าสูงที่ต้องรักษาเอาไว้และต้องให้ความสนใจและดูแลอย่างต่อเนื่อง
Silver	3.0-3.5	ลูกค้าระดับกลางที่มีความสนใจในโปรโมชั่น/แคมเปญที่ปรับ และมีความอ่อนไหวต่อราคา
Bronze	2.5-3.0	ลูกค้าที่สลับร้านซื้อสินค้า (shop s+witchers) ซึ่งอาจถูกเลื่อนเป็นกลุ่มที่มีค่าสูงขึ้นด้วยโปรโมชั่นและการกำหนดราคา มีความเสี่ยงสูงที่จะเสียลูกค้ากลุ่มนี้ไป
Standard	2.0-2.5	ลูกค้าใหม่ที่ยังมีการติดต่ोन้อยหรือใช้จ่ายน้อย บางคนอาจมีโอกาสนเลื่อนเป็นลูกค้าที่มีค่าสูงขึ้นด้วยข้อเสนอ/ส่วนลด

ที่มา: จัดทำโดยผู้แต่ง

3.4 อัลกอริทึมการจัดกลุ่ม (Clustering Algorithms) เพื่อทำการแบ่งกลุ่มลูกค้า

ในส่วนก่อนหน้านี้ เราได้สำรวจวิธีการสองวิธีที่เป็นประโยชน์ในการจัดกลุ่มลูกค้า ในส่วนนี้ เราจะก้าวไปสู่แนวทางที่ซับซ้อนขึ้นสำหรับการแบ่งกลุ่มลูกค้าโดยใช้อัลกอริทึมการจัดกลุ่ม อัลกอริทึมการจัดกลุ่มสามารถเรียนรู้จากคุณสมบัติของข้อมูลและกำหนด (หรือทำนาย) ป้ายกำกับที่เป็นตัวเลขให้กับแต่ละจุดข้อมูล ซึ่งระบุว่าจุดข้อมูลนั้นอยู่ในกลุ่มใด อัลกอริทึมการจัดกลุ่มมีอยู่หลายแบบที่สามารถใช้ในการแบ่งกลุ่มลูกค้าได้ รวมถึง K-means, DBSCAN, และ Gaussian Mixture ซึ่งจะถูกแนะนำเฉพาะอัลกอริทึม K-Means ในหัวข้อที่ 3.4.1 เนื่องจากเป็นวิธีที่นิยมใช้เป็นอย่างมากและง่ายต่อการเข้าใจมากที่สุด

3.4.1 K-Means Algorithm

K-Means algorithm เป็นหนึ่งในอัลกอริทึมการจัดกลุ่ม (clustering) ที่ได้รับความนิยมมากที่สุดสำหรับการแบ่งกลุ่มข้อมูล โดยเฉพาะในบริบทของการจัดกลุ่มลูกค้า (customer segmentation) อัลกอริทึมนี้จะจัดกลุ่มข้อมูลที่มีคุณลักษณะคล้ายคลึงกันเข้าด้วยกัน (Wang, 2025)

ขั้นตอนการทำงานของ K-Means Algorithm อธิบายโดยสรุปดังต่อไปนี้

1. **กำหนดจำนวนกลุ่ม (K)** เริ่มต้นด้วยการกำหนดจำนวนกลุ่ม (K) ที่ต้องการให้มีในการจัดกลุ่มข้อมูล เช่น $K = 3$ หมายถึงการจัดกลุ่มข้อมูลเป็น 3 กลุ่ม
2. **สุ่มเลือกจุดศูนย์กลางเริ่มต้น (Centroids)** สุ่มเลือกจุดศูนย์กลาง (centroids) จำนวน K จุดจากข้อมูลทั้งหมด เพื่อใช้เป็นจุดศูนย์กลางของแต่ละกลุ่มในขั้นตอนแรก
3. **จัดกลุ่มข้อมูล** สำหรับแต่ละข้อมูล ให้คำนวณระยะห่างระหว่างข้อมูลนั้นกับจุดศูนย์กลางทั้ง K จุด โดยใช้สมการที่ 3.8 จากนั้นจัดกลุ่มข้อมูลให้เป็นกลุ่มที่ใกล้ที่สุดกับจุดศูนย์กลาง โดยใช้สมการที่ (3.9) ตามลำดับ

ระยะห่างระหว่างจุดข้อมูล x_i และจุดศูนย์กลาง μ_j คำนวณโดยใช้ระยะห่างแบบยุคลิด (Euclidean distance) แสดงดังสมการที่ 3.8

$$d(x_i, \mu_j) = \sqrt{\sum_{k=1}^n (x_{ik} - \mu_{jk})^2} \quad (3.8)$$

โดยที่

x_i = จุดข้อมูล i

μ_j = จุดศูนย์กลาง j

n = จำนวนข้อมูลทั้งหมด

สำหรับแต่ละจุดข้อมูล x_i ให้ทำการจัดกลุ่ม C_j โดยจัดจุดข้อมูลนั้นให้อยู่ในกลุ่มที่มีจุดศูนย์กลาง μ_j ซึ่งมีระยะทางน้อยที่สุดแสดงดังสมการที่ 3.9

$$C_j = \{x_i : \min_j d(x_i, \mu_j)\} \quad (3.9)$$

4. **คำนวณจุดศูนย์กลางใหม่** หลังจากจัดกลุ่มแล้ว ให้คำนวณจุดศูนย์กลางใหม่ของแต่ละกลุ่มโดยคำนวณค่าเฉลี่ยของข้อมูลในกลุ่มนั้น ๆ โดยใช้สมการที่ 3.10 แสดงดังต่อไปนี้

$$\mu_j = \frac{1}{|C_j|} \sum_{x_i \in C_j} x_i \quad (3.10)$$

โดยที่

$|C_j|$ = จำนวนจุดข้อมูลในกลุ่ม j

5. **ทำซ้ำจนกว่าเสถียร** ทำขั้นตอนที่ 3 และ 4 ซ้ำจนกว่าจุดศูนย์กลางจะไม่เปลี่ยนแปลงหรือเปลี่ยนแปลงน้อยมาก (Convergence) แสดงว่าการจัดกลุ่มได้เสถียรแล้ว

6. **ผลลัพธ์** เมื่ออัลกอริทึมเสร็จสมบูรณ์ ข้อมูลจะถูกจัดกลุ่มตามลักษณะที่คล้ายคลึงกันภายในแต่ละกลุ่ม

ตัวอย่าง 3.3 ตัวอย่างการจัดกลุ่มลูกค้าด้วยอัลกอริทึม K-Means อย่างง่าย

การคำนวณการจัดกลุ่มลูกค้าด้วย K-Means อย่างง่าย สามารถทำได้โดยใช้การคำนวณด้วยมือ โดยมีตัวอย่างง่ายๆ ดังนี้ สมมติว่าเรามีข้อมูลลูกค้า 5 คน โดยมีคุณลักษณะ 2 อย่าง คือ "จำนวนการซื้อ" และ "มูลค่าการซื้อเฉลี่ย" ดังนี้:

ตารางที่ 3.4 ตารางข้อมูลสำหรับตัวอย่างการจัดกลุ่มลูกค้าด้วยอัลกอริทึม K-Means

ลูกค้า	จำนวนการซื้อ	มูลค่าการซื้อเฉลี่ย
A	2	10
B	3	15
C	4	20
D	5	25
E	6	30

ที่มา: จัดทำโดยผู้แต่ง

และเราต้องการจัดกลุ่มลูกค้าเหล่านี้เป็น 2 กลุ่ม ($K=2$)

ขั้นตอนการทำ ทำตามขั้นตอนที่อธิบายในหัวข้อที่ 3.4.1 แสดงดังต่อไปนี้

1. เลือกจุดศูนย์กลางเริ่มต้น

สมมติเราเลือกจุดศูนย์กลางเริ่มต้นโดยสุ่มเป็น ลูกค้า A (2, 10) และ ลูกค้า E (6, 30)

2. คำนวณระยะห่าง (Euclidean Distance)

เราจะคำนวณระยะห่างระหว่างลูกค้ากับจุดศูนย์กลางทั้งสองใช้ระยะห่าง (Euclidean Distance) ระหว่างลูกค้ากับจุดศูนย์กลาง μ_1 (ลูกค้า A) และ μ_2 (ลูกค้า E) โดยใช้สมการที่ 3.8 ตัวอย่างการคำนวณสำหรับลูกค้า B

$$\text{ระยะห่างจากศูนย์กลาง } \mu_1: \sqrt{(3-2)^2 + (15-10)^2} = \sqrt{1+25} = \sqrt{26} \approx 5.10$$

$$\text{ระยะห่างจากศูนย์กลาง } \mu_2: \sqrt{(3-6)^2 + (15-30)^2} = \sqrt{9+225} = \sqrt{234} \approx 15.29$$

3. จัดกลุ่มข้อมูล

ลูกค้าแต่ละคนจะถูกจัดกลุ่มเข้ากับจุดศูนย์กลางที่มีระยะห่างน้อยที่สุดใช้สมการที่ 3.9

4. คำนวณจุดศูนย์กลางใหม่

คำนวณจุดศูนย์กลางใหม่โดยใช้ค่าเฉลี่ยของจำนวนการซื้อและมูลค่าการซื้อเฉลี่ยของลูกค้าที่อยู่ในแต่ละกลุ่มใช้สมการที่ 3.10

5. ทำซ้ำขั้นตอนการคำนวณระยะห่างและจัดกลุ่ม

ทำซ้ำขั้นตอนที่ 2 ถึง 4 จนกว่าจุดศูนย์กลางจะไม่เปลี่ยนแปลง

ตารางที่ 3.5 ตัวอย่างผลการคำนวณในรอบแรก (1st iteration)

ลูกค้า	จำนวนการซื้อ	มูลค่าการซื้อเฉลี่ย	ระยะห่างจาก μ_1	ระยะห่างจาก μ_2	กลุ่ม
A	2	10	0	22.36	1
B	3	15	5.10	15.29	1
C	4	20	10.20	10.20	1
D	5	25	15.29	5.10	2
E	6	30	22.36	0	2

ที่มา: จัดทำโดยผู้แต่ง

6. ผลลัพธ์ ที่ได้จะแบ่งกลุ่ม 2 กลุ่มได้แก่

กลุ่มที่ 1: ลูกค้า A, B, C

กลุ่มที่ 2: ลูกค้า D, E

โดยจะมีการทำซ้ำขั้นตอนการคำนวณจุดศูนย์กลางใหม่จนกว่าจุดศูนย์กลางจะไม่เปลี่ยนแปลง จากนั้นเราจะได้ผลลัพธ์การจัดกลุ่มลูกค้าที่เสถียรแล้ว

3.5 การกำหนดนโยบายตามการแบ่งกลุ่มลูกค้า

เมื่อมีการแบ่งกลุ่มลูกค้าแล้วจำเป็นที่จะต้องมีการกำหนดนโยบายให้เหมาะสมกับแต่ละกลุ่มลูกค้าโดยมีรายละเอียดโดยสรุปดังต่อไปนี้

3.5.1 การกำหนดนโยบายสำหรับกลุ่มลูกค้าแบ่งตาม Cohort Analysis

1. **ลูกค้ากลุ่มแรกเริ่ม (Initial Cohorts)** สำหรับลูกค้าที่เพิ่งเข้ามาใหม่ (เช่น ในช่วง 3 เดือนแรก) ธุรกิจควรเสนอโบนัสหรือข้อเสนอพิเศษเพื่อกระตุ้นให้เกิดการซื้อซ้ำ (Pereira et al., 2025) เช่น คุปองส่วนลดในการซื้อครั้งที่สองหรือโปรโมชั่นซื้อหนึ่งแถมหนึ่ง เช่น หากธุรกิจสังเกตว่าลูกค้าที่ซื้อในเดือนแรกมีการซื้อต่อเนื่องในเดือนที่สอง แต่ลดลงในเดือนที่สาม อาจเสนอบัตรกำนัลส่วนลดสำหรับการซื้อในเดือนที่สามเพื่อกระตุ้นการซื้อซ้ำ

2. **ลูกค้ากลุ่มที่มีการซื้อซ้ำต่ำ (Low Retention Cohorts)** สำหรับกลุ่มลูกค้าที่มีการซื้อซ้ำต่ำ ผู้ประกอบธุรกิจสามารถใช้กลยุทธ์ retargeting หรือการส่งข้อความเฉพาะบุคคล (personalized messaging) เพื่อนำลูกค้ากลับมาซื้ออีกครั้ง เช่น การส่งอีเมลเตือนพร้อมข้อเสนอพิเศษ เช่น หากพบว่า cohort ที่ซื้อในเดือนกรกฎาคมมีอัตราการซื้อซ้ำต่ำ ผู้ประกอบธุรกิจอาจส่งอีเมลส่วนลดพิเศษสำหรับการซื้อสินค้าใหม่ในเดือนสิงหาคม

3. **ลูกค้ากลุ่มที่มีการซื้อซ้ำสูง (High Retention Cohorts)** สำหรับลูกค้าที่มีการซื้อซ้ำสูง ควรเน้นการเสริมสร้างความสัมพันธ์ระยะยาวด้วยการเสนอโปรแกรมสมาชิกที่มอบสิทธิประโยชน์พิเศษ เช่น ส่วนลดสมาชิก การเข้าถึงสินค้าก่อนใคร หรือการสะสมแต้มเพื่อแลกกับของรางวัล เช่น หาก Cohort ในเดือนมกราคมมีการซื้อซ้ำอย่างต่อเนื่อง ผู้ประกอบธุรกิจอาจเสนอโบนัสพิเศษให้สำหรับลูกค้าที่ซื้อในเดือนกุมภาพันธ์ เพื่อกระตุ้นให้มีการซื้อเพิ่มขึ้น

4. **ลูกค้ากลุ่มที่มียอดใช้จ่ายสูง (High Monetary Cohorts)** เสนอโปรโมชั่นเฉพาะที่ออกแบบมาเพื่อเพิ่มมูลค่าการซื้อเฉลี่ย (Fader et al., 2005) เช่น การจัดชุดสินค้าพิเศษ (bundle offers) หรือการเสนอสินค้าที่เกี่ยวข้องในราคาพิเศษ เช่น หากพบว่าลูกค้ากลุ่มที่ซื้อในช่วงเทศกาลมีการใช้จ่ายสูง ธุรกิจอาจเสนอชุดสินค้าพิเศษที่จับคู่กับสินค้ายอดนิยมในช่วงเทศกาลเพื่อเพิ่มยอดขาย

3.5.2 การกำหนดนโยบายสำหรับกลุ่มลูกค้าแบ่งตาม RFM (Husnah et al., 2023)

1. **ลูกค้ากลุ่มที่มีค่า RFM สูง (VIP Customers)** เสนอโปรแกรมสมาชิกพิเศษ (Fader et al., 2005) หรือโปรแกรมสะสมแต้มที่มอบสิทธิประโยชน์มากขึ้น เช่น ส่วนลดเฉพาะสมาชิก บริการส่งสินค้าฟรี หรือการเข้าถึงสินค้าก่อนใคร

2. **ลูกค้ากลุ่มที่มี Recency สูง (New Active Customers)** ใช้กลยุทธ์ในการสร้างความสัมพันธ์ระยะยาว เช่น ส่งอีเมลขอบคุณพร้อมคูปองส่วนลดสำหรับการซื้อครั้งถัดไป หรือการจัดส่งสินค้าตัวอย่างฟรีเพื่อกระตุ้นให้กลับมาซื้อซ้ำ

3. **ลูกค้ากลุ่มที่มี Frequency สูง แต่ Monetary ต่ำ (Frequent Buyers)** ส่งเสริมให้ลูกค้าซื้อสินค้าที่มีมูลค่าสูงขึ้น (Fader et al., 2005) เช่น การเสนอโปรโมชั่นซื้อ 1 แถม 1 หรือการจัดชุดสินค้าพิเศษที่มีมูลค่าสูง

4. **ลูกค้ากลุ่มที่มี RFM ต่ำ (Inactive Customers)** ใช้กลยุทธ์ในการดึงลูกค้ากลับมา (Pereira et al., 2025) เช่น การส่งโปรโมชั่นพิเศษผ่านอีเมลหรือ SMS การเสนอส่วนลดสำหรับการซื้อครั้งถัดไป หรือการส่งของขวัญพิเศษในวันเกิด

สรุป

ในบทที่แล้วได้นำเสนอแนวคิดของซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล (data-driven supply chain) ในบทนี้นำเสนอแนวคิดที่เป็นส่วนหนึ่งของบทที่แล้ว เพราะข้อมูลในซัพพลายเชนนี้นั้นมีได้หลายประเภท ซึ่งเนื้อหาในบทนี้ได้กล่าวถึงข้อมูลที่สำคัญมากได้แก่ข้อมูลของลูกค้า ซึ่งการวิเคราะห์ข้อมูลลูกค้ามีความสำคัญอย่างมากสำหรับธุรกิจ เนื่องจากข้อมูลลูกค้าช่วยให้ธุรกิจสามารถตัดสินใจได้อย่างมีข้อมูลและสร้างกลยุทธ์ที่มีประสิทธิภาพเป็นเครื่องมือสำคัญที่ช่วยให้ธุรกิจสามารถพัฒนาตนเองและสร้างความได้เปรียบในการแข่งขันได้อย่างยั่งยืน โดยเมื่อกล่าวถึงลูกค้าก็สามารถเกิดขึ้นในจุดใดก็ได้ตามในซัพพลายเชนเมื่อมีการซื้อขายเกิดขึ้น แต่ลูกค้าที่มีความสำคัญที่สุดกลับมีจุดเดียวคือ ลูกค้าที่จุดสุดท้ายหรือที่เรียกว่า ผู้บริโภคสุดท้าย (end consumers) ซึ่งจะเป็นผู้กำหนดทุกสิ่งทุกอย่างในซัพพลายเชน ตั้งแต่ การวางกลยุทธ์ทางการตลาด การออกแบบผลิตภัณฑ์ การออกแบบโครงข่ายระบบโลจิสติกส์ และอื่นๆ ประกอบกับในปัจจุบันธุรกิจต้องเผชิญตลาดที่แข่งขันอย่างรุนแรง การเข้าใจความต้องการที่แท้จริงเพื่อสร้างความพึงพอใจให้กับลูกค้าของตนเพื่อแข่งขันกับคู่แข่ง รวมถึงประสบการณ์ของลูกค้าที่มีความสำคัญเพิ่มขึ้นอย่างทวีคูณทำให้มีความจำเป็นที่ซัพพลายเชนจะต้องมุ่งเน้นลูกค้าเป็นศูนย์กลาง (customer-centric supply chain) หรือคือซัพพลายเชนที่เน้นสิ่งที่ลูกค้าต้องการจริง ๆ แล้วจึงจัดสรรเชิงกลยุทธ์ ทั้งทรัพยากร การประสานงาน และการพัฒนาความสามารถเพื่อที่จะบรรลุความคาดหวังของลูกค้า ไม่ว่าจะเป็นในอุตสาหกรรมใด ขนาดใด ภูมิภาคใด ใหม่หรือเก่า ธุรกิจทุกแห่งควรให้ความสำคัญกับลูกค้า โดยเก็บความคาดหวังของลูกค้าไว้ในใจกลางของธุรกิจและการดำเนินงานประจำวัน เพื่อที่จะเปลี่ยนซัพพลายเชนดั้งเดิมให้เป็นซัพพลายเชนที่มีลูกค้าเป็นจุดศูนย์กลางจำเป็นที่ต้องเข้าใจลูกค้าอย่างแท้จริง เครื่องมือสำคัญที่จะทำให้ความเข้าใจลูกค้าอย่างแท้จริงได้นั้นจะต้องทำการศึกษาพฤติกรรมของลูกค้า ด้วยการแบ่งกลุ่มลูกค้าโดยใช้ การวิเคราะห์ cohort การวิเคราะห์ RFM และการใช้อัลกอริทึมการจัดกลุ่มที่เป็นส่วนหนึ่งของการเรียนรู้ของเครื่อง โดยที่การแบ่งกลุ่มลูกค้าช่วยให้ธุรกิจสามารถออกแบบแคมเปญการตลาดเฉพาะสำหรับแต่ละกลุ่มลูกค้าได้ เช่น การส่งโปรโมชั่นพิเศษไปยังกลุ่มลูกค้าที่ซื้อบ่อยแต่มีมูลค่าการซื้อต่ำ เพื่อกระตุ้นให้พวกเขาซื้อสินค้าที่มีมูลค่าสูงขึ้น จึงเป็นวิธีที่มีประสิทธิภาพในการทำความเข้าใจพฤติกรรมลูกค้าและปรับปรุงกลยุทธ์การตลาดเพื่อเพิ่มประสิทธิภาพในการทำงาน

แบบฝึกหัด

- อธิบายแนวคิดของการจัดกลุ่มลูกค้า (customer segmentation) และวัตถุประสงค์หลักของการจัดกลุ่มลูกค้าในกระบวนการบริหารจัดการลูกค้า
- จาก customer cohort (x \$50) แสดงดังตารางข้างล่าง จงหา CLV (customer lifetime value) ถ้ากำหนดให้ gross margin = 55% นักศึกษาคิดว่า CAC ควรเป็นเท่าไรถึงจะเหมาะสม จงระดมสมองและ อธิบายการนำ CLV ไปใช้ประโยชน์

Month	Start Month	0	1	2	3	4	5	6	7	8	9	10	11
0	Jan	60											
1	Feb	60	21										
2	Mar	33	12	35									
3	Apr	44	20	34	40								
4	May	46	13	35	34	26							
5	Jun	46	11	34	34	25	20						
6	Jul	39	12	23	31	14	12	22					
7	Aug	30	11	30	37	18	14	15	25				
8	Sep	51	12	20	31	18	17	17	19	19			
9	Oct	58	17	31	40	22	16	15	21	19	23		
10	Nov	31	20	27	40	18	19	18	15	16	21	12	
11	Dec	56	18	22	40	14	12	11	16	17	20	12	14

- ให้นักศึกษา ทำการวิเคราะห์ RFM ของไฟล์ dataset “Retail_init_cleaned.xlsx”
 - หลังจากทำการวิเคราะห์จึงทำการคัดเลือก แบ่งกลุ่มลูกค้าตาม label แสดงดังตารางที่ 3.3 {diamond, platinum, gold, silver, bronze, standard} โดยให้นักศึกษากำหนดช่วงคะแนนในการแบ่งกลุ่มที่เหมาะสมเอง

- 3.2. ให้ทำลงในไฟล์ Microsoft Excel และนักศึกษาสามารถส่งไฟล์ดังกล่าวมายัง E-mail (natapat.ar@ssru.ac.th) หลังจากทำงานเสร็จเรียบร้อยแล้ว หมายเหตุ: ข้อมูลที่ใช้ในการวิเคราะห์ (ในไฟล์แนบ)
4. บริษัทใช้โมเดล RFM (recency, frequency, monetary) ในการแบ่งกลุ่มลูกค้า ลูกค้าคนหนึ่งมีคะแนน RFM ดังนี้: recency: 2 (1 = ล่าสุดที่สุด, 5 = ล่าสุดนานที่สุด), frequency: 1 (1 = ซื้อบ่อยที่สุด, 5 = ซื้อบ่อยน้อยที่สุด), monetary: 3 (1 = ใช้จ่ายสูงที่สุด, 5 = ใช้จ่ายน้อยที่สุด) กลยุทธ์ใดต่อไปนี้จะเหมาะสมที่สุดสำหรับกลุ่มลูกค้ากลุ่มนี้?
5. ถ้าสมมติว่าคุณใช้อัลกอริธึม K-Means เพื่อแบ่งกลุ่มลูกค้าตามพฤติกรรมการซื้อของพวกเขา หลังจากรันอัลกอริธึมแล้ว คุณพบว่ามีคลัสเตอร์หนึ่งที่มีความถี่ในการซื้อเฉลี่ยสูงกว่าคลัสเตอร์อื่น ๆ อย่างมีนัยสำคัญ แต่มีมูลค่าการทำธุรกรรมเฉลี่ยต่ำกว่าเมื่อเทียบกับคลัสเตอร์อื่น ๆ กลยุทธ์ทางการตลาดที่เป็นไปได้สำหรับกลุ่มลูกค้ากลุ่มนี้คืออะไร?
6. ในการวิเคราะห์กลุ่มลูกค้า (cohort analysis) ธุรกิจได้ระบุว่าลูกค้าที่ได้มาจากการจัดโปรโมชั่นพิเศษมีอัตราการเลิกใช้บริการสูงกว่าลูกค้าที่ได้มาจากช่องทางปกติ การดำเนินการใดที่เหมาะสมที่สุดในการแก้ไขปัญหา?
7. บริษัท B ต้องการแบ่งกลุ่มลูกค้าตามพฤติกรรมการซื้อ โดยใช้ข้อมูลดังนี้:

ลูกค้า	Recency (วัน)	Frequency (ครั้ง)	Monetary (บาท)
1	20	10	2000
2	50	5	5000
3	10	15	1000
4	60	2	8000
5	30	8	3000

- 7.1 จงใช้วิธีการ Normalization เพื่อปรับค่าของ recency, frequency, และ monetary ให้อยู่ในช่วงเดียวกัน (0-1)
- 7.2 ใช้อัลกอริธึม K-Means โดยสมมติว่า $k = 2$ เพื่อแบ่งลูกค้าออกเป็น 2 กลุ่มแสดงผลลัพธ์ว่าแต่ละลูกค้าถูกจัดอยู่ในกลุ่มใด และอธิบายลักษณะของกลุ่มนั้น ๆ
8. วิจัยข้อดีและข้อเสียของการใช้การจัดกลุ่มลูกค้าแบบดั้งเดิมกับการใช้เทคโนโลยีการวิเคราะห์ข้อมูลสมัยใหม่ เช่น machine learning ในการจัดกลุ่มลูกค้า

เอกสารอ้างอิง

- Ascarza, E., Neslin, S. A., Netzer, O., Anderson, Z., Fader, P. S., Gupta, S., Hardie, B. G. S., Lemmens, A., Libai, B., Neal, D., Provost, F., & Schrift, R. (2018). *In pursuit of enhanced customer retention management: Review, key issues, and future directions. Customer Needs and Solutions*, 5(1-2), 65-81. <https://doi.org/10.1007/s40547-017-0080-0>
- Baldi, B., Confente, I., Russo, I., & Gaudenzi, B. (2024). Consumer-centric supply chain management: A literature review, framework, and research agenda. *Journal of Business Logistics*, 45(4), 1-39. <https://doi.org/10.1111/jbl.12399>
- Chopra, S. (2019). *Supply chain management: Strategy, planning, and operation* (7th ed.). Pearson.
- Christopher, M. (2016). *Logistics and supply chain management* (5th ed.). Pearson.
- Esper, T. L., Ellinger, A. E., Stank, T. P., Flint, D. J., & Moon, M. (2020). Everything old is new again: The age of consumer-centric supply chain management. *Journal of Business Logistics*, 41(4), 286-293. <https://doi.org/10.1111/jbl.12267>.
- Fader, P. S. (2020). *Customer centricity: Focus on the right customers for strategic advantage*. Wharton School Press.
- Fedushko, S., & Ustyianovych, T. (2022). *E-commerce customers behaviour research using cohort analysis: A case study of COVID-19. Journal of Open Innovation: Technology, Market, and Complexity*, 8(1), Article 12. <https://doi.org/10.3390/joitmc8010012>.
- Husnah, M., et al. (2023). Customer segmentation analysis using LRFM-based product and brand loyalty with fuzzy C-means algorithm. In *2023 2nd International Conference for Innovation in Technology (INOCON)*. <https://doi.org/10.1109/INOCON57975.2023.10100974>
- Kumar, N. (2025). Intelligent customer segmentation: unveiling consumer patterns with machine learning. *J. Umm Al-Qura Univ. Eng.Archit.* 16, 774-783 <https://doi.org/10.1007/s43995-025-00180-7>.

- Lambert, D. M., & Cooper, M. C. (2000). Issues in supply chain management. *Industrial Marketing Management*, 29(1), 65–83. [https://doi.org/10.1016/S0019-8501\(99\)00113-3](https://doi.org/10.1016/S0019-8501(99)00113-3)
- Lemon, K. N., & Verhoef, P. C. (2016). Understanding customer experience throughout the customer journey. *Journal of Marketing*, 80(6), 69-96. <https://doi.org/10.1509/jm.15.0420>.
- Mentzer, J. T., DeWitt, W., Keebler, J. S., Min, S., Nix, N. W., Smith, C. D., & Zacharia, Z. G. (2001). Defining supply chain management. *Journal of Business Logistics*, 22(2), 1–25. <https://doi.org/10.1002/j.2158-1592.2001.tb00001.x>
- Ngai, E. W. T., Xiu, L., & Chau, D. C. K. (2009). Application of data mining techniques in customer relationship management: A literature review and classification. *Expert Systems with Applications*, 36(2), 2592–2602. <https://doi.org/10.1016/j.eswa.2008.02.021>
- Pereira, M. S., et al. (2025). Factors of customer loyalty and retention in the digital environment: A bibliometric study. *Journal of Theoretical and Applied Electronic Commerce Research*, 20(2). <https://doi.org/10.3390/jtaer20020071>
- Rahman, A. (2024). *Consumer-centric supply chain models for enhancing responsiveness and flexibility in the retail sector*. SSRN. <https://doi.org/10.2139/ssrn.5392421>.
- Wang, S., Sun, L., & Yu, Y. (2024). A dynamic customer segmentation approach by combining LRFMS and multivariate time series clustering. *Scientific Reports*, 14, 17491. <https://doi.org/10.1038/s41598-024-68621-2>.
- Wang, G. (2025). Customer segmentation in the digital marketing using a Q-learning based differential evolution algorithm integrated with K-means clustering. *PLOS ONE* 20(2): e0318519. <https://doi.org/10.1371/journal.pone.0318519>.
- Wedel, M., & Kannan, P. K. (2016). Marketing analytics for data-rich environments. *Journal of Marketing*, 80(6), 97-121. <https://doi.org/10.1509/jm.15.0413>

บทที่ 4

การวิเคราะห์ข้อมูลจัดซื้อจัดจ้าง

Procurement Data Analytics

หัวข้อ

4.1 การจัดซื้อจัดจ้าง

4.2 การคัดเลือกซัพพลายเออร์

4.3 การประเมินซัพพลายเออร์

4.4 การจัดการความสัมพันธ์กับซัพพลายเออร์

4.5 การบริหารความเสี่ยง

การจัดซื้อจัดจ้าง (procurement) เป็นกระบวนการหลักในการบริหารจัดการซัพพลายเชน ที่ครอบคลุมการจัดหาทรัพยากรต่าง ๆ ไม่ว่าจะเป็นวัตถุดิบ สินค้า หรือบริการ เพื่อให้ธุรกิจสามารถดำเนินงานได้อย่างราบรื่น กระบวนการนี้มีผลกระทบต่อต้นทุนการผลิต ทำให้การผลิตไม่สะดวก และส่งเสริมความสัมพันธ์ที่ดีกับซัพพลายเออร์ ซึ่งการจัดซื้อที่มีประสิทธิภาพจะประกอบไปด้วยขั้นตอนสำคัญได้แก่ (1) การวางแผนการจัดซื้อ (procurement planning) เป็นขั้นตอนเริ่มต้นที่สำคัญในการจัดหาทรัพยากรที่จำเป็นสำหรับองค์กร โดยกระบวนการนี้จะช่วยให้เราสามารถจัดการและควบคุมต้นทุนได้อย่างมีประสิทธิภาพ รวมถึงมั่นใจได้ว่าจะได้รับวัสดุหรือบริการที่ตรงตามความต้องการและคุณภาพที่กำหนดไว้ (2) การค้นหาและคัดเลือกซัพพลายเออร์ (supplier selection) เป็นหัวใจสำคัญในการสร้างความสัมพันธ์กับคู่ค้าทางธุรกิจที่แข็งแกร่งและมั่นคง โดยองค์กรจำเป็นต้องค้นหาและคัดเลือกซัพพลายเออร์ที่มีคุณสมบัติเหมาะสมและตอบโจทยความต้องการทางธุรกิจได้อย่างครอบคลุม (Monczka et al., 2020) (3) การบริหารความเสี่ยง (risk management) เนื่องจากความเสี่ยงในซัพพลายเชนนี้นั้นมีความซับซ้อนและหลากหลายรูปแบบ เช่น ความล่าช้าในการจัดส่งสินค้า ปัญหาด้านคุณภาพของสินค้า หรือความเสี่ยงทางการเงินของซัพพลายเออร์ ดังนั้นจึงจำเป็นต้องมีการวางแผนบริหารจัดการความเสี่ยงอย่างเป็นระบบ เพื่อลดผลกระทบที่อาจเกิดขึ้นต่อธุรกิจ (Chopra,

2019) (4) การติดตามและประเมินประสิทธิภาพของซัพพลายเออร์ (supplier performance management) หลังจากการจัดซื้อ จะต้องมีการติดตามและประเมินประสิทธิภาพของซัพพลายเออร์อย่างต่อเนื่อง เพื่อให้มั่นใจว่าซัพพลายเออร์สามารถส่งมอบสินค้าหรือบริการตามมาตรฐานที่ได้ตกลงกันไว้ (5) การบริหารจัดการความสัมพันธ์กับซัพพลายเออร์ (supplier relationship management, SRM) การบริหารความสัมพันธ์กับซัพพลายเออร์เป็นหัวใจสำคัญของกระบวนการจัดซื้อ ที่มุ่งเน้นการสร้างเชื่อมโยงอันแข็งแกร่งกับคู่ค้าทางธุรกิจ เพื่อร่วมกันพัฒนาและยกระดับกระบวนการทำงานให้มีประสิทธิภาพยิ่งขึ้น (Christopher, 2016)

4.1 การจัดซื้อจัดจ้าง

การจัดซื้อจัดจ้างนับเป็นหัวใจสำคัญของการบริหารจัดการซัพพลายเชน ช่วยให้องค์กรดำเนินงานได้อย่างมีกำไรและเป็นไปตามหลักจริยธรรม ตามที่สถาบัน Chartered Institute of Procurement and Supply (CIPS) ได้นิยามไว้ว่า การจัดซื้อจัดจ้างคือ “การซื้อสินค้าหรือบริการที่ช่วยให้องค์กรสามารถดำเนินงานได้อย่างมีกำไรและมีจริยธรรม โดยครอบคลุมกระบวนการตั้งแต่การจัดหาวัตถุดิบและบริการ ไปจนถึงการบริหารสัญญาและความสัมพันธ์กับผู้จัดหาสินค้า”

CIPS ระบุว่า การจัดซื้อจัดจ้างมีบทบาทสำคัญอย่างยิ่งต่อความสำเร็จของธุรกิจ โดยอาจคิดเป็นสัดส่วนสูงถึง 70% ของรายได้ทั้งหมด ดังนั้น การลดต้นทุนแม้เพียงเล็กน้อย ก็สามารถส่งผลกระทบต่อผลกำไรขององค์กรได้อย่างมหาศาล จากข้อมูลข้างต้น เราสามารถนิยามการจัดซื้อจัดจ้างได้อย่างกว้าง ๆ ว่าเป็นกระบวนการคัดเลือกและประเมินซัพพลายเออร์ การเจรจาต่อรองและทำสัญญา รวมถึงการบริหารจัดการความสัมพันธ์กับซัพพลายเออร์ เพื่อให้ได้มาซึ่งสินค้าหรือบริการที่มีคุณภาพในราคาที่เหมาะสม และส่งมอบตรงเวลา เพื่อสนับสนุนการดำเนินงานของธุรกิจ (Lysons & Farrington, 2020; Kundu et. al., 2020; Althabatah et. al., 2023)

คำสองคำที่มักสับสนกันคือ "การจัดซื้อจัดจ้าง" และ "การสั่งซื้อ" แม้ว่าบางครั้งจะใช้แทนกันได้ แต่ความหมายที่แท้จริงนั้นแตกต่างกันอย่างสิ้นเชิง การสั่งซื้อเป็นเพียงส่วนหนึ่งของกระบวนการจัดซื้อจัดจ้างที่กว้างกว่า โดยหมายถึงขั้นตอนที่เกี่ยวข้องโดยตรงกับการวางคำสั่งซื้อและการรับสินค้าหรือบริการจากผู้ขาย ในขณะที่การจัดซื้อจัดจ้างครอบคลุมกระบวนการทั้งหมดที่เริ่มตั้งแต่การระบุความต้องการ การค้นหาซัพพลายเออร์ การเปรียบเทียบราคา การเจรจาต่อรอง การทำสัญญา ไปจนถึงการจัดส่ง การรับสินค้า และการชำระเงิน จึงสามารถสรุปได้ว่าการสั่งซื้อเป็นเพียงส่วนย่อยของการจัดซื้อจัดจ้าง ซึ่งเป็นกระบวนการที่ซับซ้อนและครอบคลุมมากกว่า โดยการจัดซื้อจัดจ้างสามารถ

ใช้สองกลยุทธ์ที่แตกต่างกันได้แก่ การรวมกิจการแนวดิ่ง (vertical integration, VI) และ การจ้างหน่วยงานภายนอก (outsourcing, OS) (Slack et al., 2022) ซึ่งทั้งสองกลยุธินั้นมีข้อดีข้อเสียที่เหมาะสมกับสถานการณ์และเป้าหมายทางธุรกิจที่ต่างกัน โดยการทำให้ VI มีแนวคิดพื้นฐานที่บริษัทจะเข้ามาควบคุมทุกขั้นตอนในซัพพลายเชนตั้งแต่ต้นน้ำถึงปลายน้ำ หมายถึงการครอบคลุมตั้งแต่ การจัดหาวัตถุดิบ การผลิต การจัดเก็บ และการจัดจำหน่ายสินค้า เพื่อให้บริษัทสามารถบริหารจัดการและควบคุมทุกขั้นตอนได้อย่างสมบูรณ์ ซึ่งแตกต่างกับกลยุทธ์ OS ที่จะใช้บุคคลหรือองค์กรภายนอกเข้ามาทำงานบางส่วนแทนตนเองเพื่อให้บริษัทสามารถทุ่มเทความสนใจไปที่งานหลักของธุรกิจได้อย่างเต็มที่ ในขณะที่เดียวกันก็สามารถเข้าถึงความเชี่ยวชาญเฉพาะทางจากผู้เชี่ยวชาญภายนอกในงานที่ไม่ใช่แกนหลัก (core competencies) ของธุรกิจได้อย่างมีประสิทธิภาพ (Heizer et al., 2020) โดยข้อแตกต่างของทั้งสองกลยุทธ์สามารถเปรียบเทียบดังแสดงในตารางที่ 4.1

ตารางที่ 4.1 ข้อเปรียบเทียบระหว่าง Vertical Integration และ Outsourcing

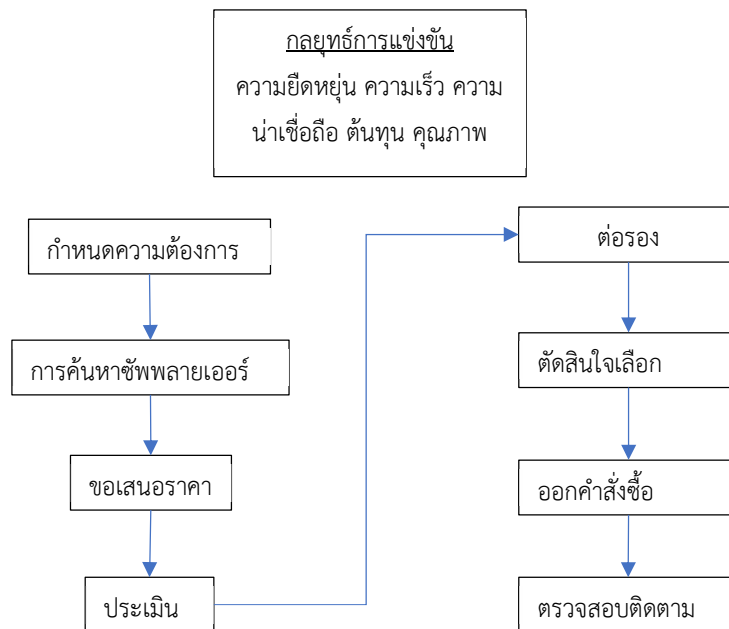
ปัจจัย	Vertical Integration	Outsourcing
การควบคุม	การควบคุมเต็มรูปแบบในกระบวนการทั้งหมด	การควบคุมลดลง เนื่องจากงานส่งออกไปภายนอก
ต้นทุน	ต้นทุนสูงในระยะเริ่มต้น แต่สามารถประหยัดในระยะยาว	ต้นทุนต่ำในระยะสั้น ไม่ต้องลงทุนในโครงสร้างพื้นฐาน
ความยืดหยุ่น	ความยืดหยุ่นน้อย เนื่องจากต้องดูแลทุกกระบวนการเอง	ความยืดหยุ่นสูง สามารถเปลี่ยนผู้ให้บริการภายนอกได้ง่าย
ความเชี่ยวชาญ	ต้องมีความเชี่ยวชาญในหลายด้านภายในองค์กร	ใช้ความเชี่ยวชาญจากผู้ให้บริการภายนอก
ความเสี่ยง	ความเสี่ยงสูงในการจัดการทุกกระบวนการ	ความเสี่ยงด้านการพึ่งพาภายนอก และการรักษาความปลอดภัยข้อมูล
การสร้างมูลค่า	สามารถสร้างมูลค่าเพิ่มได้จากการควบคุมคุณภาพและกระบวนการทั้งหมด	มุ่งเน้นที่กิจกรรมหลัก และส่งงานรองให้กับผู้เชี่ยวชาญภายนอก

ที่มา: จัดทำโดยผู้แต่ง

ดังนั้นการตัดสินใจเลือกใช้ การรวมกิจการในแนวตั้ง หรือ การจ้างหน่วยงานภายนอก ขึ้นอยู่กับ กลยุทธ์ระยะยาว ของบริษัท ทรัพยากรที่มีอยู่ และ ความสามารถในการบริหารจัดการ หากบริษัทต้องการ ควบคุมทุกขั้นตอนการผลิต และสร้าง ความได้เปรียบในการแข่งขันในระยะยาว การรวมกิจการในแนวตั้งอาจเป็นทางเลือกที่เหมาะสม แต่หากบริษัทต้องการ ความยืดหยุ่นในการดำเนินงาน และ ลดต้นทุน การจ้างหน่วยงานภายนอกอาจเป็นตัวเลือกที่ดีกว่า

4.2 การคัดเลือกซัพพลายเออร์

การคัดเลือกซัพพลายเออร์นับเป็นขั้นตอนเบื้องต้นที่สำคัญยิ่งในกระบวนการจัดซื้อจัดจ้าง เนื่องจากส่งผลโดยตรงต่อคุณภาพ ความน่าเชื่อถือ ความยืดหยุ่น และความมั่นคงในการจัดหาสินค้าหรือบริการ ซึ่งเป็นปัจจัยสำคัญในการขับเคลื่อนธุรกิจให้บรรลุเป้าหมาย การสร้างระบบคัดเลือกซัพพลายเออร์ที่รัดกุมและมีประสิทธิภาพจึงเป็นเสมือนการลงทุนที่คุ้มค่า เพราะจะช่วยให้องค์กรรักษาความสามารถในการแข่งขันในระยะยาว การคัดเลือกซัพพลายเออร์ควรสอดคล้องกับกลยุทธ์ทางธุรกิจขององค์กร (Lysons & Farrington, 2020) ตัวอย่างเช่น ในอุตสาหกรรมโทรคมนาคม บริษัทที่เน้นนวัตกรรมจะให้ความสำคัญกับซัพพลายเออร์ที่มีเทคโนโลยีล้ำสมัย ในขณะที่บริษัทที่เน้นต้นทุนจะให้ความสำคัญกับซัพพลายเออร์ที่มีราคาที่แข่งขันได้ ดังนั้น การกำหนดหลักเกณฑ์ในการคัดเลือกซัพพลายเออร์จึงต้องสอดคล้องกับเป้าหมายทางธุรกิจที่แตกต่างกันไปของแต่ละองค์กร ดังภาพที่ 4.1 แสดงให้เห็น กระบวนการคัดเลือกซัพพลายเออร์ต้องอาศัยการวิเคราะห์ที่รอบคอบ โดยผู้จัดการฝ่ายจัดซื้อต้องพิจารณาถึงกลยุทธ์ทางธุรกิจขององค์กร และเลือกซัพพลายเออร์ที่สามารถสนับสนุนให้บรรลุเป้าหมายทางธุรกิจได้อย่างมีประสิทธิภาพ



ภาพที่ 4.1 กระบวนการคัดเลือกซัพพลายเออร์

ที่มา: ภาพโดยผู้แต่ง

1. **การกำหนดความต้องการ (Requirements Definition)** เป็นการกำหนดความต้องการที่ชัดเจนและครอบคลุม ของบริษัทต่อวัตถุดิบหรือบริการ โดยพิจารณาจากปัจจัยหลักได้แก่ ลักษณะของสินค้าหรือบริการเช่นประเภทสินค้า คุณสมบัติเฉพาะ และปริมาณที่ต้องการ คุณภาพที่ต้องได้ตามมาตรฐานที่กำหนดโดยต้องผ่านเกณฑ์การตรวจสอบคุณภาพที่กำหนดไว้ ราคาตามงบประมาณที่ตั้งไว้และปฏิบัติตามเงื่อนไขการชำระเงินต่างๆ ระยะเวลาส่งมอบ และระยะเวลาในการผลิตหรือให้บริการที่จะต้องอยู่ภายในช่วงเวลาที่กำหนดไว้ นอกจากนี้ควรคำนึงถึงปัจจัยเสริมอื่น ๆ ที่สำคัญได้แก่ มาตรฐานสากลที่จะต้องมีการปฏิบัติตาม มาตรฐาน ISO GMP หรือมาตรฐานอื่น ๆ ที่เกี่ยวข้อง ความยั่งยืนเช่น การคำนึงถึงผลกระทบต่อสิ่งแวดล้อมและการใช้พลังงานอย่างมีประสิทธิภาพ รวมไปถึงการบริการหลังการขายและความน่าเชื่อถือของผู้จัดจำหน่าย เป็นต้น

การกำหนดความต้องการที่ครอบคลุมจะช่วยให้บริษัทสามารถมีข้อได้เปรียบหลายประการได้แก่ สามารถเลือกผู้จัดจำหน่ายหรือผู้ให้บริการที่เหมาะสมโดยเปรียบเทียบ

ข้อเสนอจากหลาย ๆ ราย สามารถควบคุมคุณภาพของสินค้าหรือบริการโดยตรวจสอบให้ได้ตามมาตรฐานที่กำหนดไว้ ช่วยลดต้นทุนโดยการเลือกผู้จัดจำหน่ายที่มีราคาที่เหมาะสมและมีเงื่อนไขการชำระเงินที่เอื้ออำนวย ช่วยสร้างความสัมพันธ์ที่ดีกับผู้จัดจำหน่ายโดยการสื่อสารความต้องการอย่างชัดเจนและสร้างความเข้าใจร่วมกัน

2. **การค้นหาซัพพลายเออร์ (Supplier Sourcing)** เพื่อจัดหาวัตถุดิบหรือบริการที่ตรงตามความต้องการของบริษัท เป็นขั้นตอนสำคัญที่ส่งผลต่อประสิทธิภาพและความสำเร็จของธุรกิจ (Baily et al., 2015) โดยทั่วไปแล้วมีช่องทางหลักในการค้นหาซัพพลายเออร์ได้หลายช่องทาง เช่น ค้นหาจากฐานข้อมูลซัพพลายเออร์ที่มีอยู่ โดยทั่วไปบริษัทมักจะมีฐานข้อมูลซัพพลายเออร์ที่เคยติดต่อหรือใช้บริการอยู่แล้ว การกลับไปตรวจสอบฐานข้อมูลนี้เป็นจุดเริ่มต้นที่ดีในการค้นหาซัพพลายเออร์ที่อาจตอบโจทย์ความต้องการในปัจจุบันได้ ค้นหาตามงานแสดงสินค้าโดยเข้าร่วมงานแสดงสินค้าที่เกี่ยวข้องกับอุตสาหกรรมของบริษัทจะเป็นโอกาสที่ดีในการพบปะกับซัพพลายเออร์จำนวนมากในที่เดียว และสามารถเปรียบเทียบสินค้าหรือบริการของแต่ละรายได้อย่างใกล้ชิด ค้นหาจากแพลตฟอร์มออนไลน์ ในปัจจุบันมีแพลตฟอร์มออนไลน์มากมายที่ให้บริการเชื่อมต่อระหว่างผู้ซื้อและผู้ขาย เช่น เว็บไซต์ B2B, Marketplace ต่างๆ การใช้แพลตฟอร์มออนไลน์ช่วยให้การค้นหาซัพพลายเออร์เป็นไปอย่างรวดเร็วและสะดวกสบายมากขึ้น หรือค้นหาจากการแนะนำจากผู้จัดหาสินค้ารายอื่น โดยการขอคำแนะนำจากซัพพลายเออร์รายอื่นที่เคยทำงานด้วย หรือจากคู่ค้าทางธุรกิจ เป็นอีกหนึ่งช่องทางที่น่าสนใจ เพราะข้อมูลจากบุคคลที่สามจะช่วยให้ตัดสินใจเลือกซัพพลายเออร์ได้อย่างมั่นใจมากขึ้น อย่างไรก็ตามการตัดสินใจเลือกใช้ช่องทางใดยังขึ้นอยู่กับปัจจัยหลายประการ เช่น ประเภทของวัตถุดิบหรือบริการที่ต้องการ ความเร่งด่วนของการจัดหางบประมาณที่ตั้งไว้และขนาดของบริษัท

3. **การออกคำร้องขอเสนอราคา (Request for Quotation, RFQ) หรือข้อเสนอ (Request for Proposal - RFP)** เป็นการที่บริษัทส่งคำขอให้ซัพพลายเออร์ที่มีศักยภาพเสนอราคา (RFQ) หรือส่งข้อเสนอ (RFP) ที่รวมถึงรายละเอียดเกี่ยวกับสินค้า บริการ เงื่อนไขการจัดส่ง และราคา ซัพพลายเออร์ต้องตอบสนองต่อคำขอนี้ด้วยข้อมูลเกี่ยวกับสินค้าและบริการที่สามารถจัดหาได้ รวมถึงราคาและเงื่อนไขการชำระเงิน

4. การประเมินซัพพลายเออร์ (Supplier Evaluation) คือการเลือกซัพพลายเออร์ที่ดีที่สุดเป็นขั้นตอนสำคัญในการดำเนินธุรกิจ การประเมินซัพพลายเออร์จึงเป็นกระบวนการที่จำเป็นเพื่อให้ได้มาซึ่งซัพพลายเออร์ที่มีคุณภาพและเหมาะสมกับความต้องการขององค์กร โดยมีปัจจัยหลักกำหนดตามความต้องการที่ระบุไว้ในข้อที่ 1 เช่น การประเมินจากคุณภาพของวัตถุดิบหรือบริการที่ซัพพลายเออร์เสนอมา เพื่อให้มั่นใจว่าสินค้าหรือบริการที่ได้รับตรงตามมาตรฐานที่กำหนดไว้ และตอบสนองความต้องการของลูกค้าได้อย่างมีประสิทธิภาพ การประเมินจากต้นทุน นอกจากคุณภาพแล้ว ต้นทุนก็เป็นปัจจัยสำคัญที่ต้องพิจารณาในการเลือกซัพพลายเออร์ การเปรียบเทียบราคาที่ซัพพลายเออร์เสนอมา รวมถึงต้นทุนที่เกี่ยวข้องอื่นๆ เช่น ค่าขนส่ง ค่าภาษี เป็นต้น จะช่วยให้สามารถเลือกซัพพลายเออร์ที่ให้ราคาที่เหมาะสมและคุ้มค่า การประเมินจากความสามารถในการส่งมอบสินค้าหรือบริการตรงตามเวลาที่กำหนดไว้เป็นสิ่งสำคัญอย่างยิ่ง เพราะจะส่งผลต่อกระบวนการผลิตและการส่งมอบสินค้าให้กับลูกค้า จึงทำให้โดยทั่วไปจะต้องประเมินความน่าเชื่อถือและประสิทธิภาพในการจัดส่งของซัพพลายเออร์ด้วย การประเมินถัดไปคือ การประเมินจากความมั่นคงทางการเงิน เป็นการวิเคราะห์ความมั่นคงทางการเงินของซัพพลายเออร์เป็นสิ่งจำเป็น เพื่อลดความเสี่ยงที่อาจเกิดขึ้น เช่น การไม่สามารถส่งมอบสินค้าได้ตามกำหนด หรือการล้มละลายของซัพพลายเออร์ การประเมินอันดับสุดท้ายได้แก่ การประเมินจากประวัติการทำงาน โดยจะทำการพิจารณาประวัติการทำงานของซัพพลายเออร์ เช่น ประสบการณ์ในการทำงานร่วมกับบริษัทอื่น ๆ หรือข้อมูลรีวิวกจากลูกค้ารายอื่น จะช่วยให้สามารถประเมินความน่าเชื่อถือและความสามารถของซัพพลายเออร์ได้อย่างรอบด้าน (Ho et al., 2010)

5. การเจรจาต่อรอง (Negotiation)

เมื่อคัดเลือกซัพพลายเออร์ที่ผ่านเกณฑ์การประเมินแล้ว ขั้นตอนต่อไปคือการเจรจาต่อรองเพื่อหาข้อตกลงที่เป็นประโยชน์แก่ทั้งสองฝ่าย โดยประเด็นสำคัญที่ต้องเจรจา ได้แก่ ราคา สินค้าและบริการที่จะต้องมีการตกลงราคาที่เหมาะสม ซึ่งสอดคล้องกับคุณภาพและปริมาณที่ต้องการ การเจรจาเกี่ยวกับเงื่อนไขการชำระเงิน โดยจะต้องมีการกำหนดระยะเวลาและวิธีการชำระเงินที่สะดวกทั้งผู้ซื้อและผู้ขาย การเจรจาเกี่ยวกับเงื่อนไขการส่งมอบ ต้องทำการตกลงเรื่องเวลาการส่งมอบสินค้าหรือบริการ และวิธีการขนส่ง การเจรจาเกี่ยวกับการรับประกัน โดยจะต้องกำหนดเงื่อนไขการรับประกันสินค้าหรือบริการ เพื่อความมั่นใจใน

คุณภาพ และการเจรจาเกี่ยวกับเงื่อนไขสัญญา โดยจะต้องมีการทำสัญญาที่ครอบคลุม รายละเอียดต่าง ๆ เช่น ระยะเวลาของสัญญา ข้อกำหนดในการปรับเปลี่ยนราคา และบริการ หลังการขาย นอกจากนี้ในขั้นตอนการเจรจาต่อรองนั้นทุกฝ่ายของผู้เจรจาย่อมต้องมีเป้าหมาย สูงสุดของการเจรจา คือการสร้างผลประโยชน์ร่วมกัน (win-win) (van Weele, 2018) เพื่อให้ทั้งบริษัทและซัพพลายเออร์ได้รับประโยชน์สูงสุด และสร้างความสัมพันธ์ทางธุรกิจที่ยั่งยืนในระยะยาว

6. การตัดสินใจคัดเลือก (Final Supplier Selection)

หลังจากการเจรจาต่อรองเสร็จสิ้น บริษัทจะทำการตัดสินใจเลือกซัพพลายเออร์ที่ดีที่สุด โดยพิจารณาจากผลลัพธ์ของการประเมินและการเจรจาต่อรอง โดยที่ซัพพลายเออร์ที่ได้รับการคัดเลือกจะได้รับสัญญาหรือข้อตกลงที่ตกลงร่วมกัน

7. การออกคำสั่งซื้อ (Purchase Order - PO)

หลังจากเลือกซัพพลายเออร์ บริษัทจะออกคำสั่งซื้อ (PO) ให้กับซัพพลายเออร์ที่ได้รับการคัดเลือก โดย PO จะระบุรายละเอียดเกี่ยวกับจำนวนสินค้า ราคาที่ตกลงไว้ เงื่อนไขการชำระเงิน และเวลาการส่งมอบ

8. การตรวจสอบและติดตามผลการดำเนินงาน (Supplier Performance Monitoring)

หลังจากเริ่มต้นการทำงานกับซัพพลายเออร์ที่ได้รับการคัดเลือกแล้ว ต้องมีการติดตามผลการดำเนินงานเพื่อให้มั่นใจว่าซัพพลายเออร์สามารถจัดส่งสินค้าหรือบริการตามเงื่อนไขที่กำหนดไว้ โดยที่การประเมินผลการดำเนินงานของซัพพลายเออร์ควรใช้ตัวชี้วัด (KPIs) เช่น ความตรงต่อเวลาในการส่งมอบสินค้า คุณภาพของสินค้า และการตอบสนองต่อปัญหาหรือคำร้องเรียน ประกอบการประเมิน (Krause et al., 1998)

4.3 การประเมินซัพพลายเออร์

ในส่วนนี้ขอกล่าวถึงการประเมินซัพพลายเออร์โดยละเอียดเนื่องจากเป็นขั้นตอนที่จำเป็นอย่างยิ่งในการบริหารจัดการซัพพลายเชนและกระบวนการจัดซื้อ (Leenders et al., 2006) โดยมีวัตถุประสงค์หลักเพื่อคัดเลือกและทำงานร่วมกับซัพพลายเออร์ที่สามารถตอบสนองความต้องการขององค์กรได้อย่างครอบคลุม ทั้งในแง่ของคุณภาพสินค้าหรือบริการ ตรงตามกำหนดเวลา และมี

ประสิทธิภาพสูงสุด โดยมีกิจกรรมและปัจจัยหลักในการประเมินดังนี้ (Alkahtani et. al., 2019; Hou et al., 2022)

1. **คุณภาพของสินค้าและบริการ (Quality)** การประเมินคุณภาพของสินค้าหรือบริการที่จัดซื้อจากซัพพลายเออร์เป็นขั้นตอนสำคัญที่ส่งผลโดยตรงต่อคุณภาพของผลิตภัณฑ์หรือบริการที่องค์กรนำเสนอสู่ตลาด ตัวชี้วัดที่บ่งบอกถึงคุณภาพของสินค้าหรือบริการที่ได้รับจากซัพพลายเออร์ ได้แก่ อัตราสินค้าที่มีข้อบกพร่อง จำนวนการคืนสินค้า และผลการตรวจสอบคุณภาพที่ไม่ผ่านเกณฑ์มาตรฐานที่กำหนด (Beamon, 1999)
2. **การส่งมอบตรงเวลา (On-time Delivery)** ซัพพลายเออร์จำเป็นต้องส่งมอบสินค้าตรงตามกำหนดเวลาที่ตกลงกันไว้ เพื่อให้กระบวนการผลิตและการดำเนินงานขององค์กรเป็นไปอย่างราบรื่นและต่อเนื่องตัวชี้วัดสำคัญที่ใช้ในการประเมิน ได้แก่ สัดส่วนของการส่งมอบสินค้าที่ตรงตามกำหนดเวลาที่กำหนดไว้ และจำนวนครั้งที่เกิดความล่าช้าซึ่งเป็นตัวชี้วัดที่บ่งบอกถึงความสำเร็จในการส่งมอบสินค้า และช่วยให้สามารถระบุปัญหาและหาแนวทางแก้ไขได้อย่างตรงจุด
3. **ต้นทุนและราคาที่แข่งขันได้ (Cost and Price Competitiveness)** โดยราคาที่ซัพพลายเออร์เสนอมานั้น ต้องสอดคล้องกับคุณภาพของสินค้าหรือบริการที่ได้รับ นอกจากนี้ราคายังต้องมีความสามารถในการแข่งขันได้ในตลาดด้วย เพื่อให้ธุรกิจของเราสามารถรักษาความได้เปรียบและดึงดูดลูกค้า ในการพิจารณาเลือกซัพพลายเออร์ เราควรประเมินต้นทุนรวม (Total Cost of Ownership - TCO) ซึ่งหมายถึงไม่เพียงแต่ราคาซื้อ แต่รวมถึง ค่าใช้จ่ายอื่น ๆ ที่เกี่ยวข้องทั้งหมด เช่น ค่าขนส่ง ค่าเก็บรักษาสินค้า ค่าบริการหลังการขาย และอื่น ๆ ที่อาจเกิดขึ้นในระยะยาว การพิจารณา TCO จะช่วยให้เราเห็นภาพรวมของค่าใช้จ่ายที่แท้จริง และเลือกซัพพลายเออร์ที่คุ้มค่าที่สุด
4. **ความสามารถทางเทคนิคและนวัตกรรม (Technical Capability and Innovation)** ซัพพลายเออร์จะต้องมีความสามารถทางเทคนิคสูง เพื่อผลิตสินค้าให้ตรงตามข้อกำหนดและมีคุณภาพ รวมถึงสามารถนำนวัตกรรมมาประยุกต์ใช้ในการพัฒนาสินค้าอย่างต่อเนื่องความสามารถในการพัฒนาผลิตภัณฑ์ใหม่และปรับปรุงกระบวนการผลิตอย่างสม่ำเสมอ ถือเป็นตัวชี้วัดสำคัญที่บ่งบอกถึงความสามารถทางเทคนิคและนวัตกรรมของซัพพลายเออร์

5. **ความมั่นคงทางการเงิน (Financial Stability)** ความแข็งแกร่งทางการเงินของซัพพลายเออร์เป็นปัจจัยสำคัญที่ส่งผลต่อความต่อเนื่องในการจัดหาสินค้าและบริการในระยะยาว การประเมินความเสี่ยงทางการเงินของซัพพลายเออร์จึงเป็นขั้นตอนที่จำเป็นในการบริหารจัดการซัพพลายเชนอย่างมีประสิทธิภาพ การวิเคราะห์งบการเงิน การตรวจสอบประวัติเครดิต และการประเมินศักยภาพในการชำระหนี้ เป็นเครื่องมือที่ช่วยให้เราสามารถประเมินความมั่นคงทางการเงินของซัพพลายเออร์ได้อย่างแม่นยำ
6. **ความยืดหยุ่นและการตอบสนอง (Flexibility and Responsiveness)** ซัพพลายเออร์ที่มีความยืดหยุ่นสูงในการปรับตัวและตอบสนองต่อความต้องการที่เปลี่ยนแปลงของลูกค้านั้นเป็นสินทรัพย์ที่สำคัญยิ่งต่อองค์กร โดยเฉพาะอย่างยิ่งในสถานการณ์ที่ความต้องการของตลาดมีความผันผวนสูง เช่น การเพิ่มหรือลดจำนวนคำสั่งซื้ออย่างฉับพลัน หรือการเปลี่ยนแปลงระยะเวลาในการส่งมอบ ความยืดหยุ่นของซัพพลายเออร์ช่วยให้องค์กรสามารถตอบสนองต่อความต้องการของตลาดได้อย่างรวดเร็ว ลดความเสี่ยงจากการขาดแคลนสินค้า และสร้างความพึงพอใจให้กับลูกค้า
7. **ความรับผิดชอบต่อสังคมและสิ่งแวดล้อม (Corporate Social Responsibility - CSR)** ซัพพลายเออร์ทุกแห่งควรให้ความสำคัญกับการดำเนินธุรกิจอย่างมีจริยธรรม รับผิดชอบต่อสังคม และคำนึงถึงผลกระทบต่อสิ่งแวดล้อม การประเมินและตรวจสอบในด้านนี้จะช่วยสร้างความเชื่อมั่นให้กับองค์กร และส่งเสริมภาพลักษณ์ที่ดีต่อสาธารณะ ((Humphreys et al., 2003; Chaitongrat & Areerakulkan, 2021) การประเมินที่ครอบคลุมควรเน้นประเด็นสำคัญ ได้แก่ การประเมินด้านมาตรฐานจริยธรรม โดยตรวจสอบว่าซัพพลายเออร์ปฏิบัติตามหลักจริยธรรมทางธุรกิจอย่างเคร่งครัดหรือไม่? การประเมินความรับผิดชอบต่อสังคมโดยประเมินการมีส่วนร่วมของซัพพลายเออร์ในการพัฒนาชุมชน และการส่งเสริมคุณภาพชีวิตของผู้คน การประเมินด้านสิ่งแวดล้อมโดยตรวจสอบการปฏิบัติตามกฎหมายและมาตรฐานด้านสิ่งแวดล้อม รวมถึงการจัดการขยะ การลดการปล่อยมลพิษ และการใช้ทรัพยากรอย่างยั่งยืน การประเมินด้านแรงงานโดยประเมินเงื่อนไขการทำงานของพนักงานของซัพพลายเออร์ว่าเป็นไปตามกฎหมายแรงงานหรือไม่ เช่น ค่าจ้าง สวัสดิการ และความปลอดภัยในการทำงาน

การดำเนินการตามมาตรฐานเหล่านี้ จะช่วยให้องค์กรมีภาพลักษณ์ที่ดี สร้างความเชื่อมั่นให้กับลูกค้า พนักงาน นักลงทุน และผู้มีส่วนได้ส่วนเสียทุกฝ่าย ช่วยลดความเสี่ยงเพื่อป้องกันปัญหาทางกฎหมายและความเสียหายต่อชื่อเสียง และสร้างความยั่งยืนโดยสนับสนุนการพัฒนาที่ยั่งยืนทั้งในด้านเศรษฐกิจ สังคม และสิ่งแวดล้อม

8. การบริหารความสัมพันธ์กับซัพพลายเออร์ (Supplier Relationship Management, SRM) การสร้างความสัมพันธ์อันดีกับซัพพลายเออร์ ไม่เพียงแต่ช่วยให้การดำเนินงานเป็นไปอย่างราบรื่น แต่ยังส่งผลให้เกิดการทำงานร่วมกันที่มีประสิทธิภาพสูงขึ้น โดยมีตัวชี้วัดที่สำคัญ ได้แก่ การสื่อสารที่เปิดเผย ความโปร่งใสในการดำเนินงาน และการร่วมมือกันแก้ไขปัญหาและพัฒนาโครงการต่างๆ

4.3.1 วิธีการประเมินซัพพลายเออร์ด้วยการตัดสินใจแบบลำดับชั้นเชิงวิเคราะห์

(Analytic Hierarchy Process, AHP)

AHP หรือกระบวนการวิเคราะห์ลำดับชั้น เป็นเครื่องมือที่ใช้ในการตัดสินใจเชิงปริมาณที่พัฒนาขึ้นโดย Thomas L. Saaty (Saaty & Vargas, 2012) วิธีการนี้ช่วยให้เราสามารถจัดการกับปัญหาการตัดสินใจที่ซับซ้อนที่มีหลายปัจจัย โดยอาศัยหลักการเปรียบเทียบแบบคู่ (Pairwise Comparison) เพื่อหาค่าน้ำหนักสัมพัทธ์ของแต่ละปัจจัย จากนั้นจึงนำค่าน้ำหนักที่ได้มาคำนวณหาค่าความสำคัญรวมของแต่ละทางเลือก ทำให้เราสามารถเลือกทางเลือกที่ดีที่สุดได้อย่างเป็นระบบ (Ho, 2008) ขั้นตอนของวิธี AHP แสดงดังต่อไปนี้

1. การเปรียบเทียบเชิงคู่ (Pairwise Comparison) (Bruno et al., 2012)

ขั้นตอนแรกของกระบวนการวิเคราะห์ลำดับชั้น (AHP) คือ การสร้างเมทริกซ์เปรียบเทียบเชิงคู่สำหรับแต่ละเกณฑ์ เพื่อประเมินความสำคัญสัมพัทธ์ระหว่างทางเลือกต่างๆ ในแต่ละเกณฑ์นั้น ๆ โดยพิจารณาจากความชอบของผู้ตัดสินใจ ในตัวอย่างเกณฑ์ราคา เราจะนำค่าที่แสดงในตารางที่ 4.2 มาใช้ในการเปรียบเทียบเชิงคู่ระหว่างระบบ X Y และ Z โดยกำหนดให้ P_{ij} แทนระดับความชอบที่ทางเลือก i มีต่อทางเลือก j บนเกณฑ์ราคา (1) ตัวอย่างเช่น หากผู้ตัดสินใจชอบราคาของระบบ X มากกว่าระบบ Y อย่างมาก ค่า P_{XY} จะเท่ากับ 7 (2) อีกตัวอย่างหนึ่งเช่น หากผู้ตัดสินใจชอบราคาของระบบ X มากกว่าระบบ Z อย่างมาก และชอบราคาของระบบ Y มากกว่าระบบ Z ในระดับปานกลาง ค่า P_{XZ} จะเท่ากับ 7 และ P_{YZ} จะเท่ากับ 3 ตามลำดับ

ค่าการเปรียบเทียบเชิงคู่เหล่านี้จะถูกนำมาสร้างเป็นเมทริกซ์เปรียบเทียบเชิงคู่ ดังแสดงในตารางที่ 4.3 ซึ่งจะนำไปสู่การคำนวณค่าความสำคัญสัมพัทธ์ของแต่ละทางเลือกในขั้นตอนต่อไป

ตารางที่ 4.2 สเกลการเปรียบเทียบเชิงคู่ใน AHP

คะแนน	ความหมาย	ความสำคัญสัมพัทธ์
1	เท่ากัน	ความสำคัญเท่ากัน
3	พอสมควร	ความสำคัญพอสมควร
5	มากกว่า	ความสำคัญมากกว่า
7	มากที่สุด	ความสำคัญมากที่สุด
9	มากที่สุดในระดับสูงสุด	ความสำคัญสูงสุด
2, 4, 6, 8	ค่าเชิงกลางหรือค่าระหว่างค่า	ความสำคัญในระดับกลาง

ที่มา: ดัดแปลงมาจาก (Saaty, 2016)

ตัวอย่าง 4.1 การให้คะแนนเปรียบเทียบเชิงคู่

คะแนนการเปรียบเทียบเชิงคู่ระหว่างทางเลือกซัพพลายเออร์ X Y และ Z ในเกณฑ์เรื่องราคาแสดงดังตารางที่ 4.3

ตารางที่ 4.3 เมตริกเปรียบเทียบเชิงคู่

	X	Y	Z
X	1.000	5.000	7.000
Y	0.200	1.000	3.000
Z	0.143	0.333	1.000
ผลรวม	1.343	6.333	11.000

หมายเหตุ: ผลรวม ($1.343 = P_{xx} + P_{yx} + P_{zx} = 1.000 + 0.200 + 0.143$)

ที่มา: จัดทำโดยผู้แต่ง

ข้อสังเกต: ค่าตามเส้นทแยงมุมของเมทริกซ์จะมีค่าเท่ากับ 1 ($P_{xx} = P_{yy} = P_{zz} = 1$) ทั้งหมด เพราะเมื่อเปรียบเทียบราคาของระบบเดียวกันควรที่จะมีค่าเท่ากัน และค่าคะแนนของ $P_{ij} = \frac{1}{P_{ji}}$ เช่น ค่าคะแนน $P_{yx} = \frac{1}{P_{xy}} = \frac{1}{5} = 0.2$

4.3.2 การแปลงคะแนนเปรียบเทียบเป็นคะแนนมาตรฐาน (Normalizing Comparison)

ขั้นตอนต่อจากการให้คะแนนคือการทำให้คะแนนเปรียบเทียบเป็นคะแนนปกติมาตรฐาน โดยใช้สูตรแสดงดังสมการที่ 4.1 และ 4.2 กำหนดให้เมทริกซ์คะแนนมีขนาดเท่ากับ $m \times n$ (m = จำนวนแถวทั้งหมด และ n = จำนวนหลักทั้งหมด)

$$x_{norm,i,j} = \frac{x_{ij}}{\sum_{i=1}^m x_{ij}} \text{ สำหรับ } j = 1, \dots, n \quad (4.1)$$

$$S_i = \frac{\sum_{j=1}^n x_{norm,i,j}}{n} \text{ สำหรับ } i = 1, \dots, m \quad (4.2)$$

โดยที่ S_i คือ ค่าคะแนนน้ำหนักของแต่ละเกณฑ์หรือทางเลือก

$x_{norm,i,j}$ คือ ค่าคะแนนมาตรฐาน

ตัวอย่าง 4.2 การแปลงคะแนนเปรียบเทียบเป็นคะแนนมาตรฐาน

สำหรับการแปลงคะแนนเปรียบเทียบเป็นคะแนนมาตรฐานสามารถทำได้โดยใช้สมการที่ 4.1 คือการนำค่าคะแนนแต่ละตัวมาหารด้วยผลรวม ผลลัพธ์แสดงดังตารางที่ 4.4

ตารางที่ 4.4 การแปลงคะแนนเปรียบเทียบเป็นคะแนนมาตรฐาน

	X	Y	Z	คะแนน (S_i)
X	0.745	0.789	0.636	0.724
Y	0.149	0.158	0.273	0.193
Z	0.106	0.053	0.091	0.083

หมายเหตุ: $x_{norm,1,1} = x_{norm,X,X} = 1.000/1.343=0.745$

ที่มา: จัดทำโดยผู้แต่ง

4.3.3 ความสอดคล้อง (Consistency)

ความสอดคล้องคือ ความสอดคล้องของคะแนนเปรียบเทียบเชิงคู่ที่ให้ความสอดคล้องในเชิงตรรกะโดยไม่มีการขัดแย้งในลำดับความสำคัญเช่น

สมมติว่าเรามีตัวเลือกอยู่ 3 ตัวเลือก ได้แก่ A, B, และ C ตามลำดับ ถ้าเราให้คะแนนเปรียบเทียบเชิงคู่ระหว่างทางเลือก A และทางเลือก B ว่า A สำคัญกว่า B เป็น 3 เท่า หลังจากนั้นจึงทำการเปรียบเทียบระหว่าง B กับ C พบว่า B สำคัญกว่า C เป็น 4 เท่า ดังนั้นเมื่อเปรียบเทียบระหว่าง A กับ C ถ้าเราให้คะแนน A สำคัญ น้อยกว่า หรือ เท่ากับ C ก็แสดงว่าการให้คะแนนนั้นไม่สมเหตุสมผล หรืออีกกรณี เราให้ A สำคัญกว่า C จริง แต่สำคัญกว่าไม่มากพอก็อาจส่งผลให้เกิดความไม่สอดคล้องเช่นเดียวกัน

ในกรณีของวิธี AHP ได้มีการสร้างสมการที่ใช้ตรวจสอบความสอดคล้องของคะแนนแสดงดังสมการที่ 4.3 และ 4.4

$$Consistency\ Index\ (CI) = \frac{\lambda - n}{n - 1} \quad (4.3)$$

$$Consistency\ Ratio\ (CR) = \frac{CI}{RI} \quad (4.4)$$

โดยที่ λ = ค่าเฉลี่ยความสอดคล้องทั้งหมด

n = จำนวนทางเลือก

RI = ดัชนีความเหมาะสมแสดงดังตารางที่ 4.5

ตารางที่ 4.5 ค่า RI สำหรับวิธี AHP

n	2	3	4	5	6	7	8
RI	0	0.58	0.9	1.12	1.24	1.32	1.41

ที่มา: ดัดแปลงมาจาก (Saaty, 2016)

ค่า CR ที่คำนวณได้ควรมีค่าน้อยกว่า 0.1 ($CR \leq 0.10$) ถึงจะมีความสอดคล้อง ถ้าคะแนน CR มากกว่า 0.10 แสดงว่าคะแนนประเมินที่ให้นั้นไม่มีความสอดคล้องและจำเป็นที่จะต้องมีการปรับคะแนนใหม่จนกว่าจะน้อยกว่า 0.10

ตัวอย่าง 4.3 การคำนวณหาค่าดัชนีความสอดคล้อง (CR)

การคำนวณหาค่าดัชนีความสอดคล้องต้องเริ่มจากการหาค่าคะแนนความสอดคล้องของแต่ละแถว โดยนำเวกเตอร์คะแนน (S_i) ไปคูณกับค่าคะแนนการเปรียบเทียบของแต่ละแถว (ตารางที่ 4.3) และหารด้วยค่าคะแนนตำแหน่งนั้น เช่น

$$\text{Consistency measure ของ } X = \frac{0.724 \times 1 + 0.193 \times 5 + 0.083 \times 7}{0.724} = 3.141$$

$$\text{Consistency measure ของ } Y = \frac{0.724 \times 0.200 + 0.193 \times 1 + 0.083 \times 3}{0.193} = 3.043$$

$$\text{Consistency measure ของ } Z = \frac{0.724 \times 0.143 + 0.333 \times 1 + 1.000 \times 3}{0.193} = 3.014$$

ตารางที่ 4.6 ค่าคะแนนความสอดคล้อง

	X	Y	Z	ค่าน้ำหนัก (S_i)	คะแนน ความ สอดคล้อง
X	0.745	0.789	0.636	0.724	3.141
Y	0.149	0.158	0.273	0.193	3.043
Z	0.106	0.053	0.091	0.083	3.014

ที่มา: จัดทำโดยผู้แต่ง

$$\lambda = \frac{(3.141 + 3.043 + 3.014)}{3} = 3.066$$

$$CI = \frac{\lambda - n}{n - 1} = \frac{3.066 - 3}{3 - 1} = 0.033, CR = \frac{0.033}{0.58} = 0.057$$

ตัวอย่าง 4.4 การใช้วิธี AHP ประเมินซัพพลายเออร์ X Y และ Z

สมมติว่าบริษัทต้องการประเมินซัพพลายเออร์ X Y และ Z โดยใช้เกณฑ์การประเมินได้แก่ ราคา คุณภาพ ความรวดเร็วในการจัดส่ง ผู้ประเมินทำการประเมินคะแนนความสำคัญของเกณฑ์การประเมินโดยวิธีเปรียบเทียบเชิงคู่แสดงดังตารางที่ 4.7 ค่าคะแนนเปรียบเทียบเชิงคู่ของซัพพลายเออร์ตามราคา ค่าคะแนนเปรียบเทียบเชิงคู่ของซัพพลายเออร์ตามคุณภาพ และ ค่าคะแนนเปรียบเทียบเชิงคู่ของซัพพลายเออร์ตามความเร็วในการส่ง แสดงดังตารางที่ 4.8- 4.10 ตามลำดับ

ตารางที่ 4.7 ค่าคะแนนเปรียบเทียบเชิงคู่เกณฑ์การประเมิน

	ราคา	คุณภาพ	ความเร็ว
ราคา	1.000	0.333	0.250
คุณภาพ	3.000	1.000	0.500
ความเร็ว	4.000	2.000	1.000
ผลรวม	8.000	3.333	1.750

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 4.8 ค่าคะแนนประเมินชีพพลายเออร์ตามราคา

	X	Y	Z
X	1.000	5.000	7.000
Y	0.200	1.000	3.000
Z	0.143	0.333	1.000
ผลรวม	1.343	6.333	11.000

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 4.9 ค่าคะแนนประเมินชีพพลายเออร์ตามคุณภาพ

	X	Y	Z
X	1.000	0.333	2.000
Y	3.000	1.000	5.000
Z	0.500	0.200	1.000
ผลรวม	4.500	1.533	8.000

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 4.10 ค่าคะแนนประเมินชีพพลายเออร์ตามความเร็ว

	X	Y	Z
X	1.000	0.500	0.333
Y	2.000	1.000	0.500
Z	3.000	2.000	1.000
ผลรวม	6.000	3.500	1.833

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนต่อไปของวิธี AHP คือ การแปลงคะแนนให้เป็นค่ามาตรฐาน (Normalization Comparison) โดยใช้สมการที่ 4.2 ยกตัวอย่างเช่น

$$x_{norm, \text{ราคา, คุณภาพ}} = \frac{x_{ij}}{\sum_{i=1}^m x_{ij}} \quad \text{ถ้า } j = \text{คุณภาพ}$$

$$x_{norm, \text{ราคา, คุณภาพ}} = \frac{0.333}{3.333} = 0.100$$

ทำลักษณะเดียวกันกับทุกค่าในตารางที่ 4.7 จะได้ผลลัพธ์แสดงดังตารางที่ 4.11 ในส่วนของคะแนนมาตรฐาน

ตารางที่ 4.11 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้องของเกณฑ์

คะแนนมาตรฐาน					
	ราคา	คุณภาพ	ความเร็ว	ค่าน้ำหนัก	ค่าความสอดคล้อง
ราคา	0.125	0.100	0.143	0.123*	3.006
คุณภาพ	0.375	0.300	0.286	0.320*	3.019
ความเร็ว	0.500	0.600	0.571	0.557*	3.030
ที่มา: จัดทำโดยผู้แต่ง					CR = 0.016

ทำการแปลงคะแนนให้เป็นค่ามาตรฐานสำหรับคะแนนประเมินตามเกณฑ์แต่ละตัวแสดงดังตารางที่ 4.12-4.14 ตามลำดับ

ตารางที่ 4.12 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้อง ตามราคา

คะแนนมาตรฐาน					
	X	Y	Z	X	ค่าความสอดคล้อง
X	0.745	0.789	0.636	0.724*	3.141
Y	0.149	0.158	0.273	0.193*	3.043
Z	0.106	0.053	0.091	0.083*	3.014
ที่มา: จัดทำโดยผู้แต่ง					CR = 0.057

ตารางที่ 4.13 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้อง ตามคุณภาพ

คะแนนมาตรฐาน					
	X	Y	Z	X	ค่าความสอดคล้อง
X	0.222	0.217	0.250	0.230*	3.003
Y	0.667	0.652	0.625	0.648*	3.007
Z	0.111	0.130	0.125	0.122*	3.001
ที่มา: จัดทำโดยผู้แต่ง					CR = 0.003

ตารางที่ 4.14 ค่าคะแนนมาตรฐาน ค่าน้ำหนักและ ค่าอัตราส่วนความสอดคล้อง ตามความเร็ว

คะแนนมาตรฐาน					
	X	Y	Z	X	ค่าความสอดคล้อง
X	0.167	0.143	0.182	0.164*	3.004
Y	0.333	0.286	0.273	0.297*	3.008
Z	0.500	0.571	0.545	0.539*	3.015
ที่มา: จัดทำโดยผู้แต่ง					CR = 0.008

หลังจากนั้นเราสามารถที่จะคำนวณคะแนนถ่วงน้ำหนักของซัพพลายเออร์แต่ละรายได้โดยการนำค่าคะแนนของซัพพลายเออร์ที่ถูกประเมินตามเกณฑ์ต่างๆมาคูณกับค่าน้ำหนักของเกณฑ์แสดงดังตารางที่ 4.15

ตารางที่ 4.15 ค่าคะแนนถ่วงน้ำหนักของซัพพลายเออร์แต่ละราย

ซัพพลายเออร์				
เกณฑ์	X	Y	Z	ค่าน้ำหนัก
ราคา	0.724	0.193	0.083	0.123
คุณภาพ	0.230	0.648	0.122	0.320
ความเร็ว	0.164	0.297	0.539	0.557
คะแนนถ่วงน้ำหนัก	0.254	0.397**	0.350	1.000

ที่มา: จัดทำโดยผู้แต่ง

คะแนนถ่วงน้ำหนักสามารถคำนวณได้ดังต่อไปนี้

$$\text{คะแนนถ่วงน้ำหนักของซัพพลายเออร์ X} = 0.724 \times 0.1233 + 0.230 \times 0.320 + 0.164 \times 0.557 = 0.254$$

จากตารางที่ 4.15 เราสามารถสรุปได้ว่า ซัพพลายเออร์ Y มีคะแนนประเมินถ่วงน้ำหนักมากที่สุดมาเป็นอันดับที่หนึ่ง ตามมาด้วยซัพพลายเออร์ Z และ X ตามลำดับ ดังนั้นในการประเมินซัพพลายเออร์โดยใช้เกณฑ์ด้าน ราคา คุณภาพ และความเร็วในการส่งสินค้า พบว่า เราควรที่จะเลือกซัพพลายเออร์ Y

4.4 การจัดการความสัมพันธ์กับซัพพลายเออร์

การจัดการความสัมพันธ์กับซัพพลายเออร์เป็นกลยุทธ์สำคัญที่องค์กรนำมาใช้เพื่อสร้างและรักษาความสัมพันธ์ที่ดีกับคู่ค้าทางธุรกิจ (Monczka et al., 2020) โดยมีเป้าหมายเพื่อเพิ่มประสิทธิภาพในการดำเนินงาน ลดต้นทุน และสร้างความได้เปรียบในการแข่งขัน โดยมีหลักการสำคัญได้แก่

1. การแบ่งประเภทของซัพพลายเออร์ หนึ่งในหลักการสำคัญของ SRM คือการแบ่งประเภทของซัพพลายเออร์ เนื่องจากซัพพลายเออร์แต่ละรายมีบทบาทและความสำคัญแตกต่างกันไป การแบ่งประเภทจะช่วยให้องค์กรสามารถจัดสรรทรัพยากรและวางแผนกลยุทธ์ในการบริหารจัดการความสัมพันธ์ได้อย่างเหมาะสม เช่นการแบ่งเป็นซัพพลายเออร์หลักและรองโดย (1) ซัพพลายเออร์หลัก (strategic suppliers) เป็นซัพพลายเออร์ที่มีบทบาทสำคัญต่อธุรกิจ เช่น ซัพพลายเออร์วัตถุดิบหลักที่หายากหรือมีเทคโนโลยีเฉพาะตัว ซัพพลายเออร์ที่ร่วมพัฒนาผลิตภัณฑ์ใหม่ หรือซัพพลายเออร์ที่มีส่วนแบ่งทางการตลาดสูง องค์กรควรให้ความสำคัญกับซัพพลายเออร์ประเภทนี้เป็นพิเศษ โดยสร้างความสัมพันธ์ที่ใกล้ชิด สร้างความไว้วางใจ และร่วมกันพัฒนาธุรกิจในระยะยาว (2) ซัพพลายเออร์รอง (non-strategic suppliers) เป็นซัพพลายเออร์ที่จำหน่ายสินค้าหรือบริการทั่วไปที่หาได้ง่าย และมีผลกระทบต่อธุรกิจในระดับปานกลาง องค์กรอาจใช้กลยุทธ์ในการจัดการความสัมพันธ์ที่แตกต่างออกไป เช่น การจัดซื้อในปริมาณมากเพื่อลดต้นทุน หรือการหมุนเวียนซัพพลายเออร์เพื่อเพิ่มอำนาจในการต่อรอง

2. **การพัฒนาความร่วมมือระยะยาว** เป้าหมายสูงสุดของ SRM คือการสร้างพันธมิตรที่แข็งแกร่งกับซัพพลายเออร์ เพื่อร่วมกันสร้างสรรค์นวัตกรรมใหม่ ๆ พัฒนาคุณภาพผลิตภัณฑ์ และบริการให้ดียิ่งขึ้น ตอบสนองความต้องการของตลาดที่เปลี่ยนแปลงอย่างรวดเร็ว

3. **การสื่อสารและการแลกเปลี่ยนข้อมูล** ความโปร่งใสและการสื่อสารอย่างต่อเนื่องเป็นรากฐานสำคัญของการบริหารความสัมพันธ์กับซัพพลายเออร์ (SRM) การแลกเปลี่ยนข้อมูลที่เป็นประโยชน์และทันต่อเหตุการณ์ เช่น ข้อมูลความต้องการของตลาด สถานะสินค้าคงคลัง หรือแผนการผลิตที่กำลังดำเนินอยู่ ช่วยให้ซัพพลายเออร์สามารถปรับตัวและตอบสนองต่อความต้องการของธุรกิจได้อย่างมีประสิทธิภาพและตรงจุดยิ่งขึ้น (Bag et al., 2022)

4. **การบริหารความเสี่ยง** ความเสี่ยงจากการพึ่งพาซัพพลายเออร์มากเกินไปหรือการส่งมอบสินค้าที่ไม่ตรงตามกำหนดของซัพพลายเออร์ เป็นสิ่งที่องค์กรทุกแห่งต้องเผชิญและให้ความสำคัญในการบริหารจัดการ เพื่อลดผลกระทบที่อาจเกิดขึ้นต่อการดำเนินงานและภาพลักษณ์ขององค์กร การมีแผนสำรองที่แข็งแกร่ง จึงเป็นสิ่งจำเป็นอย่างยิ่ง หนึ่งในกลยุทธ์ที่สำคัญคือ การกระจายการจัดซื้อจากหลายแหล่ง เพื่อลดความเสี่ยงจากการพึ่งพาซัพพลายเออร์เพียงรายเดียว นอกจากนี้ การตรวจสอบสถานะทางการเงินและการดำเนินงานของซัพพลายเออร์อย่างสม่ำเสมอ ก็เป็นอีกหนึ่งแนวทางที่ช่วยให้องค์กรสามารถคาดการณ์และป้องกันปัญหาที่อาจเกิดขึ้นได้ล่วงหน้า

ตัวอย่าง 4.5 กรณีศึกษาการจัดการความสัมพันธ์กับซัพพลายเออร์ของบริษัท TOYOTA

หนึ่งในปัจจัยสำคัญที่ทำให้ Toyota ประสบความสำเร็จอย่างสูงในอุตสาหกรรมยานยนต์ คือการบริหารจัดการความสัมพันธ์กับซัพพลายเออร์ (SRM) ที่เป็นเลิศ ความสัมพันธ์อันยาวนานและแน่นแฟ้นกับคู่ค้า ทำให้ Toyota สามารถสร้างระบบการผลิตที่ราบรื่นและมีประสิทธิภาพ ซึ่งเป็นรากฐานสำคัญที่สนับสนุนการเติบโตอย่างต่อเนื่องของบริษัท โดยมีกลยุทธ์ที่สำคัญได้แก่

1. การพัฒนาความร่วมมือระยะยาว (Long-term Partnerships)

โตโยต้าแตกต่างจากบริษัทอื่นๆ ตรงที่ให้ความสำคัญกับการสร้างความสัมพันธ์ที่ยั่งยืนกับซัพพลายเออร์มากกว่าการมุ่งเน้นผลประโยชน์ระยะสั้น การทำงานร่วมกันอย่างใกล้ชิดและการสนับสนุนด้านต่างๆ ทำให้ซัพพลายเออร์สามารถพัฒนาขีดความสามารถและเติบโตไปพร้อมกับโตโยต้า ซึ่งเป็นรากฐานสำคัญที่ทำให้สามารถสร้างผลิตภัณฑ์ที่มีคุณภาพสูงและเป็นที่ยอมรับของผู้บริโภคทั่วโลก

2. โครงการปรับปรุงประสิทธิภาพอย่างต่อเนื่อง (Continuous Improvement)

โครงการปรับปรุงประสิทธิภาพอย่างต่อเนื่องของ Toyota นั้นขับเคลื่อนด้วยหลักการ Kaizen ซึ่งส่งเสริมให้ทุกภาคส่วนร่วมมือกันในการปรับปรุงอย่างไม่หยุดยั้ง ซัพพลายเออร์ของ Toyota จึงมีบทบาทสำคัญในการร่วมกันยกระดับกระบวนการผลิต เพื่อให้เกิดประสิทธิภาพสูงสุด ลดต้นทุน และลดของเสียในทุกขั้นตอนของสายการผลิต

3. การให้ความช่วยเหลือทางการเงินและการจัดการ (Support and Collaboration)

เพื่อสร้างความแข็งแกร่งให้กับซัพพลายเชน โตโยต้าให้ความสำคัญกับการสนับสนุนซัพพลายเออร์อย่างต่อเนื่อง ด้วยการนำเสนอทั้งความรู้และทรัพยากรที่จำเป็น ไม่ว่าจะเป็นการให้คำปรึกษาเชิงลึก การสนับสนุนทางการเงิน หรือโอกาสในการพัฒนาทักษะ ผ่านการฝึกอบรมและโปรแกรมต่างๆ ทำให้ซัพพลายเออร์สามารถเพิ่มประสิทธิภาพการทำงานและตอบสนองความต้องการของโตโยต้าได้อย่างมีประสิทธิภาพยิ่งขึ้น

ผลลัพธ์ที่ได้จากการจัดการ SRM ของ TOYOTA: ด้วยการพัฒนาอย่างต่อเนื่อง ซัพพลายเออร์ของ Toyota จึงสามารถผลิตชิ้นส่วนที่มีคุณภาพสูง ส่งผลให้ Toyota รักษามาตรฐานคุณภาพของรถยนต์ได้อย่างสม่ำเสมอ นอกจากนี้ยังสามารถลดความเสี่ยงจากปัจจัยต่างๆ เนื่องจากการสร้างความสัมพันธ์อันแข็งแกร่งและความร่วมมืออย่างใกล้ชิดกับซัพพลายเออร์ ช่วยให้ Toyota สามารถบริหารจัดการซัพพลายเชนได้อย่างมีประสิทธิภาพ และ สร้างกระบวนการผลิตที่มีความยืดหยุ่นทำให้สามารถปรับตัวรับมือสถานการณ์ที่เปลี่ยนแปลงได้อย่างรวดเร็ว

4.5 การบริหารความเสี่ยง (Risk Management)

การบริหารความเสี่ยงของซัพพลายเออร์ เป็นกระบวนการสำคัญที่ช่วยให้บริษัทสามารถลดความเสี่ยงที่อาจเกิดขึ้นจากการพึ่งพาผู้จัดหาสินค้า วัตถุดิบ หรือบริการในซัพพลายเชน การจัดการที่ดีจะช่วยให้บริษัทสามารถดำเนินงานได้อย่างต่อเนื่องและมีประสิทธิภาพสูงสุด รวมทั้งช่วยป้องกันปัญหาที่อาจส่งผลกระทบต่อการผลิตหรือการบริการ (Helmold et al., 2022) กระบวนการบริหารความเสี่ยงของซัพพลายเออร์อธิบายดังต่อไปนี้

1. การระบุความเสี่ยงของซัพพลายเออร์ (Risk Identification) การระบุความเสี่ยงคือ

การค้นหาปัจจัยที่อาจเป็นภัยหรือความเสี่ยงจากซัพพลายเออร์ ซึ่งอาจส่งผลกระทบต่อการทำงานของบริษัท ความเสี่ยงที่พบบ่อยประกอบด้วย ความเสี่ยงด้านการขาดแคลนวัตถุดิบเพราะซัพพลายเออร์อาจไม่สามารถจัดหาวัตถุดิบเพียงพอได้ เนื่องจากข้อจำกัดด้านทรัพยากร การผลิต หรือปัจจัยภายนอก เช่น สถานะเศรษฐกิจ เป็นต้น ความเสี่ยงจากการล่าช้าในการส่งมอบเนื่องมาจากการส่ง

มอบที่ไม่ตรงเวลาอาจเกิดจากหลายสาเหตุ เช่น ปัญหาด้านโลจิสติกส์ ภัยธรรมชาติ หรือปัญหาการดำเนินงานของซัพพลายเออร์ ซึ่งส่งผลให้บริษัทไม่สามารถผลิตหรือส่งมอบสินค้าได้ตามกำหนด ความเสี่ยงจากคุณภาพของสินค้าหรือบริการต่ำ เนื่องมาจากสินค้าหรือบริการที่ไม่เป็นไปตามมาตรฐานที่ต้องการอาจทำให้เกิดความสูญเสียในกระบวนการผลิต การต้องส่งคืนสินค้า หรือการชดเชยความเสียหาย ความเสี่ยงทางการเงินของซัพพลายเออร์เกิดจากปัญหาทางการเงินที่อาจส่งผลให้ไม่สามารถดำเนินงานต่อไปได้ หรือเกิดการล้มละลาย ซึ่งจะทำให้การจัดหาหยุดชะงัก ความเสี่ยงสุดท้ายคือ ความเสี่ยงจากปัจจัยทางภูมิศาสตร์หรือการเมือง หากซัพพลายเออร์ ตั้งอยู่ในประเทศที่มีปัญหาทางการเมือง ความไม่แน่นอนทางกฎหมาย หรือภัยธรรมชาติ บริษัทอาจเผชิญกับความเสี่ยงจากการไม่สามารถจัดส่งสินค้าได้

2. การประเมินความเสี่ยงของซัพพลายเออร์ (Risk Assessment) เมื่อระบุความเสี่ยงได้แล้ว จำเป็นต้องทำการประเมินความเสี่ยงเหล่านั้นว่ามีผลกระทบมากน้อยเพียงใด และมีโอกาสเกิดขึ้นแค่ไหน โดยการประเมินจะพิจารณาจาก ปัจจัยหลักสามประการได้แก่ (1) ระดับความรุนแรง (Severity) การพิจารณาว่าความเสี่ยงที่เกิดขึ้นจะส่งผลกระทบต่อธุรกิจในระดับใด หากเกิดขึ้น เช่น ความเสียหายทางการเงินหรือการล่าช้าในสายการผลิต (2) ความถี่ในการเกิดขึ้น (Frequency) การวิเคราะห์ว่าเหตุการณ์เหล่านั้นมีโอกาสเกิดขึ้นบ่อยเพียงใด และ (3) ความพร้อมในการตอบสนอง (Response Capability) บริษัทมีความสามารถเพียงใดในการตอบสนองหรือแก้ไขเมื่อความเสี่ยงเกิดขึ้น เช่น การมีแผนสำรองหรือการเปลี่ยนซัพพลายเออร์ อย่างรวดเร็ว

3. การบรรเทาความเสี่ยง (Risk Mitigation) หลังจากประเมินความเสี่ยงแล้ว บริษัทจะต้องหาวิธีลดความเสี่ยง โดยวิธีที่นิยมใช้เช่น การกระจายความเสี่ยงด้วยการใช้ซัพพลายเออร์หลายราย (supplier diversification) แทนที่จะพึ่งพาซัพพลายเออร์ รายเดียว ควรมีซัพพลายเออร์หลายรายจากแหล่งที่แตกต่างกัน เพื่อให้สามารถเปลี่ยนหรือเลือกใช้รายอื่นได้เมื่อมีปัญหา การสร้างความสัมพันธ์ระยะยาว (long-term relationships) ทำได้โดยการสร้างความสัมพันธ์ที่ดีและยั่งยืนกับซัพพลายเออร์ ช่วยสร้างความไว้วางใจและความมุ่งมั่นในการส่งมอบสินค้าและบริการที่มีคุณภาพอย่างสม่ำเสมอ การทำสัญญาที่มีความยืดหยุ่น (flexible contracts) โดยการกำหนดสัญญาที่มีเงื่อนไขชัดเจนในเรื่องการส่งมอบ คุณภาพสินค้า และการชดเชยในกรณีเกิดปัญหา รวมถึงความยืดหยุ่นในการปรับเปลี่ยนข้อตกลงเมื่อเกิดเหตุการณ์ที่ไม่คาดคิด หรือการทำระบบตรวจสอบผู้จัดหา (supplier audits) โดยทำการตรวจสอบประสิทธิภาพของซัพพลายเออร์ อย่างสม่ำเสมอ ทั้งในด้าน

คุณภาพ การส่งมอบ และสถานะทางการเงิน เพื่อให้มั่นใจว่าซัพพลายเออร์ ยังคงมีความสามารถในการตอบสนองต่อความต้องการของบริษัท

4. การใช้เทคโนโลยีในการติดตามและประเมินผล (Technology and Monitoring)

การใช้เทคโนโลยีและข้อมูลในการบริหารความเสี่ยงช่วยให้สามารถติดตามสถานะของซัพพลายเออร์ได้แบบเรียลไทม์และคาดการณ์ความเสี่ยงล่วงหน้า เช่น การสร้างระบบจัดการซัพพลายเชน ที่ช่วยติดตามการดำเนินงานของซัพพลายเออร์ ทั้งในด้านการส่งมอบ คุณภาพ และประสิทธิภาพการผลิต หรือการวิเคราะห์ข้อมูลเชิงลึก (predictive analytics) โดยการใช้ข้อมูลในอดีตและแบบจำลองเชิงคาดการณ์เพื่อตรวจจับสัญญาณความเสี่ยง เช่น การเปลี่ยนแปลงในสถานะทางการเงินของซัพพลายเออร์ หรือการล่าช้าของการส่งมอบ (El Baz & Ruel, 2021)

5. การวางแผนการตอบสนองและกรณีฉุกเฉิน (Response and Contingency Planning) เพื่อป้องกันความเสี่ยงที่เกิดขึ้นอย่างกะทันหัน บริษัทควรมีแผนสำรองหรือวิธีการจัดการในกรณีที่ซัพพลายเออร์ ไม่สามารถดำเนินการได้ เช่น การจัดเตรียมผู้จัดหารอง (secondary suppliers) ที่สามารถให้บริการได้ทันทีหากผู้จัดหารายหลักประสบปัญหา การเพิ่มสินค้าคงคลังสำรอง (inventory buffers) เพื่อให้มีสินค้าคงคลังสำรองในกรณีที่ปัญหาการจัดส่งจากซัพพลายเออร์ หรือ การพัฒนากลยุทธ์เชิงรุก (proactive strategies) โดยจะต้องมีการติดตามข้อมูลที่เกี่ยวข้องกับความเสี่ยงอย่างสม่ำเสมอ เช่น ข้อมูลทางการเงิน ภูมิศาสตร์ หรือเศรษฐกิจ เพื่อปรับตัวและเตรียมแผนสำรองล่วงหน้าได้ (El Baz & Ruel, 2021)

6. การสื่อสารและการทำงานร่วมกัน (Collaboration and Communication) การสื่อสารที่มีประสิทธิภาพและความโปร่งใสระหว่างบริษัทและ ซัพพลายเออร์ เป็นปัจจัยสำคัญในการลดความเสี่ยง โดยการสร้างความไว้วางใจและความโปร่งใส การสร้างความสัมพันธ์ที่เปิดเผยและโปร่งใสกับ ซัพพลายเออร์ ช่วยให้ ซัพพลายเออร์ แจ้งเตือนถึงปัญหาที่อาจเกิดขึ้นล่วงหน้า หรือการเจรจาต่อรองและทำงานร่วมกันในกรณีที่เกิดปัญหา เช่น การแก้ไขปัญหาเกี่ยวกับคุณภาพสินค้า การปรับแผนการส่งมอบ หรือตกลงเรื่องการชดเชยในกรณีที่มีความล่าช้า

7. การติดตามและประเมินผลอย่างต่อเนื่อง (Continuous Monitoring and Review)

การบริหารความเสี่ยงของซัพพลายเออร์ เป็นกระบวนการที่ต้องทำอย่างต่อเนื่อง โดยมีการติดตามประสิทธิภาพและความเสี่ยงของซัพพลายเออร์ อย่างสม่ำเสมอในด้านต่างๆ เช่น การตรวจสอบประสิทธิภาพการจัดส่ง (performance monitoring) โดยติดตามว่าผู้จัดหาสามารถส่งมอบสินค้าได้

ตรงเวลาและตามมาตรฐานที่กำหนดไว้หรือไม่ การประเมินคุณภาพสินค้าและบริการ (quality assessment) โดยตรวจสอบคุณภาพของสินค้าที่ได้รับเพื่อให้แน่ใจว่าเป็นไปตามข้อกำหนดที่ตกลงกันไว้ และการทบทวนแผนความเสี่ยง (risk review) โดยจะต้องทบทวนแผนบริหารความเสี่ยงและปรับปรุงตามสถานการณ์ใหม่ๆ ที่เกิดขึ้น

สรุป

การจัดซื้อจัดจ้างเป็นกระบวนการหลักในการบริหารจัดการซัพพลายเชน (SCM) ที่ครอบคลุมการจัดหาทรัพยากรต่างๆ ไม่ว่าจะเป็นวัตถุดิบ สินค้า หรือบริการ เพื่อให้ธุรกิจสามารถดำเนินงานได้อย่างราบรื่น กระบวนการนี้มีผลกระทบต่อต้นทุนการผลิต การผลิตไม่สะดุด และความสัมพันธ์กับซัพพลายเออร์ ซึ่งการจัดซื้อที่มีประสิทธิภาพจะประกอบไปด้วยขั้นตอนสำคัญได้แก่ (1) การวางแผนการจัดซื้อ (2) การค้นหาและคัดเลือกซัพพลายเออร์ (3) การบริหารความเสี่ยง (4) การติดตามและประเมินประสิทธิภาพของซัพพลายเออร์ (5) การบริหารจัดการความสัมพันธ์กับซัพพลายเออร์ โดยขั้นตอนที่สำคัญที่มีผลต่อการจัดซื้อจัดจ้างที่มีประสิทธิภาพได้แก่ขั้นตอนการค้นหาและคัดเลือกซัพพลายเออร์ ซึ่งเริ่มจาก กำหนดความต้องการโดยระบุสินค้าหรือบริการที่ต้องการรวมถึงคุณสมบัติและปริมาณที่จำเป็น หลังจากนั้นจึงทำการค้นหาซัพพลายเออร์โดยรวบรวมรายชื่อซัพพลายเออร์ที่ศักยภาพจากแหล่งข้อมูลต่างๆ ทำการประเมินเบื้องต้นโดยตรวจสอบความน่าเชื่อถือ ความสามารถและประสบการณ์ของซัพพลายเออร์ ส่งคำขอเสนอราคา (RFQ) หรือคำขอข้อเสนอ (RFP) ไปยังซัพพลายเออร์ที่คัดเลือก วิเคราะห์และประเมินข้อเสนอด้านราคา คุณภาพ เงื่อนไขการส่งมอบ และเงื่อนไขอื่นๆ ทำการเจรจาต่อรองเงื่อนไขที่ดีที่สุดกับซัพพลายเออร์ที่มีศักยภาพ แล้วจึงทำการเลือกซัพพลายเออร์ที่เหมาะสมที่สุดเพื่อสรุปข้อตกลงและจัดทำสัญญาอย่างเป็นทางการ เมื่อมีการจัดทำสัญญาเรียบร้อยแล้วจะต้องมีการติดตามประเมินผลประสิทธิภาพของซัพพลายเออร์เพื่อทำการปรับปรุงในอนาคต โดยในขั้นตอนของการประเมินนั้นจะทำการประเมินตามเกณฑ์ต่างๆ ไม่ว่าจะเป็น คุณภาพ ความรวดเร็วในการจัดส่ง ความมั่นคงทางการเงิน ความยืดหยุ่นและการตอบสนอง ความรับผิดชอบต่อสิ่งแวดล้อมและสังคม หรือต้นทุนและราคาที่แข่งได้ วิธีที่นิยมใช้ในการประเมินได้แก่วิธี AHP ด้วยสาเหตุที่ว่าวิธี AHP ช่วยในการจัดการกับปัญหาที่ซับซ้อนเนื่องจากมีเกณฑ์การประเมินที่เกี่ยวข้องหลายเกณฑ์ นอกจากนั้นยังสามารถประเมินได้ทั้งข้อมูลเชิงปริมาณและเชิงคุณภาพ เป็นที่ยอมรับอย่างกว้างขวางในวงวิชาการและอุตสาหกรรม และอื่น ๆ ในหัวข้อสุดท้ายได้กล่าวถึงการบริหารความเสี่ยงของซัพพลายซึ่งเป็นกระบวนการที่องค์กรใช้ในการระบุ ประเมิน และจัดการความเสี่ยงที่เกี่ยวข้องกับซัพพลายเชน เพื่อป้องกันปัญหาที่อาจส่งผลกระทบต่อการดำเนินงานและเป้าหมายทางธุรกิจ โดยมีขั้นตอนหลักคือ การระบุความเสี่ยงที่เริ่มจากการค้นหาปัจจัยที่อาจก่อให้เกิดความเสี่ยง เช่น ปัญหาการเงินของซัพพลายเออร์ ภัยธรรมชาติ หลังจากนั้นจึงทำการประเมินความเสี่ยง โดยทำการวิเคราะห์ความน่าจะเป็นและผลกระทบของความเสี่ยงที่ระบุเพื่อกำหนดระดับความสำคัญ ขั้นตอนต่อมาคือ การวางแผนจัดการความเสี่ยง โดยการกำหนดกลยุทธ์และมาตรการในการลดหลีกเลี่ยง หรือยอมรับความเสี่ยง เช่น การหาแหล่งซัพพลายเออร์สำรอง เมื่อมีแผนการแล้วจะต้องมีการดำเนินการตามแผนที่วางไว้ สุดท้ายก็คือ การติดตามผลการดำเนินการและการปรับปรุงแก้ไขปัญหาต่าง ๆ ที่เกิดขึ้น

แบบฝึกหัด

1. อธิบายขั้นตอนหลักของกระบวนการจัดซื้อจัดจ้าง (Procurement Process) ตั้งแต่เริ่มต้นจนจบ พร้อมยกตัวอย่างการประยุกต์ใช้ในธุรกิจที่คุณรู้จัก
2. ปัจจัยใดบ้างที่ควรพิจารณาในการเลือกซัพพลายเออร์ (Supplier Selection) และการพิจารณาปัจจัยเหล่านี้สำคัญต่อการจัดซื้อจัดจ้างอย่างไร โปรดอธิบาย
3. การทำสัญญาจัดซื้อจัดจ้าง (Procurement Contracts) มีประเภทใดบ้าง? อธิบายความแตกต่างระหว่างแต่ละประเภทและระบุข้อดีข้อเสียของแต่ละประเภท
4. อธิบายประเภทความเสี่ยงที่เกี่ยวข้องกับซัพพลายเออร์ (Supplier Risks) ที่บริษัทอาจเผชิญ และเสนอแนวทางการจัดการเพื่อลดความเสี่ยงเหล่านี้
5. การประเมินความเสี่ยงของซัพพลายเออร์ (Supplier Risk Assessment) มีขั้นตอนอย่างไรบ้าง? ยกตัวอย่างวิธีการประเมินที่บริษัทสามารถใช้ได้ และอธิบายว่าการประเมินนี้ช่วยให้การจัดการความเสี่ยงมีประสิทธิภาพอย่างไร
6. การกระจายความเสี่ยงผ่านการใช้ซัพพลายเออร์หลายราย (Supplier Diversification) มีประโยชน์อย่างไรในการจัดการความเสี่ยงของซัพพลายเออร์? พร้อมทั้งยกตัวอย่างกรณีศึกษาหรือเหตุการณ์จริงที่แสดงถึงความสำคัญของการกระจายความเสี่ยง
7. กระบวนการสร้างความสัมพันธ์เชิงกลยุทธ์ (Strategic Supplier Relationships) มีองค์ประกอบอะไรบ้าง? ยกตัวอย่างการใช้กลยุทธ์ความสัมพันธ์นี้ในธุรกิจใดธุรกิจหนึ่งและอธิบายว่ามีผลกระทบต่อการทำงานของบริษัทอย่างไร
8. บริษัทกำลังพิจารณาซัพพลายเออร์ 3 ราย โดยพิจารณาตามเกณฑ์ 3 ข้อคือ ราคา (Price), คุณภาพ (Quality) และการส่งมอบตรงเวลา (Delivery Timeliness) มีข้อมูลการเปรียบเทียบเชิงคู่ระหว่างเกณฑ์ดังนี้:

เกณฑ์	ราคา	คุณภาพ	การส่งมอบตรงเวลา
ราคา	1	3	5
คุณภาพ	1/3	1	2
การส่งมอบตรงเวลา	1/5	1/2	1

และคะแนนที่ให้ซัพพลายเออร์แต่ละรายแสดงดังตารางต่อไปนี้

ซัพพลายเออร์	ราคา	คุณภาพ	การส่งมอบตรงเวลา
ซัพพลายเออร์ A	7	8	6
ซัพพลายเออร์ B	6	9	8
ซัพพลายเออร์ C	8	7	7

จงตอบคำถามต่อไปนี้

- 1) สร้างตารางเปรียบเทียบเชิงคู่ (Pairwise Comparison Matrix) และคำนวณเวกเตอร์น้ำหนัก (Priority Vector) ของเกณฑ์ต่าง ๆ
- 2) ใช้เวกเตอร์น้ำหนักเพื่อคำนวณคะแนนรวมของซัพพลายเออร์แต่ละราย
- 3) วิเคราะห์และเลือกซัพพลายเออร์ที่เหมาะสมที่สุด

เอกสารอ้างอิง

- Althabatah, A., Yaqot, M., Menezes, B., & Kerbache, L. (2023). *Transformative procurement trends: Integrating Industry 4.0 technologies for enhanced procurement processes*. *Logistics*, 7(3), Article 63.
<https://doi.org/10.3390/logistics7030063>.
- Alkahtani, M., Al-Ahmari, A., Kaid, H., & Sonboa, M. (2019). Comparison and evaluation of multi-criteria supplier selection approaches: A case study. *Advances in Mechanical Engineering*, 11(2), 1-19.
<https://doi.org/10.1177/1687814018822926>.
- Bag, S., Choi, T. M., Rahman, M. S., Srivastava, G., & Singh, R. K. (2022). Examining collaborative buyer-supplier relationships and social sustainability in the "new normal" era: the moderating effects of justice and big data analytical intelligence. *Annals of operations research*, 1–46. Advance online publication.
<https://doi.org/10.1007/s10479-022-04875-1>
- Baily, P., Farmer, D., Crocker, B., Jessop, D., & Jones, D. (2015). *Procurement principles and management* (11th ed.). Pearson.
- Beamon, B. M. (1999). Measuring supply chain performance. *International Journal of Operations & Production Management*, 19(3), 275–292.
<https://doi.org/10.1108/01443579910249714>
- Chaitongrat, S., & Areerakulkan, N. (2021). Application of analytic hierarchy process method for sustainable supplier selection: A case study of ABC Company. *Journal of Logistics and Supply Chain College*, 7(2), 5-17.
- Chopra, S. (2019). *Supply chain management: Strategy, planning, and operation* (7th ed.). Pearson.
- Christopher, M. (2016). *Logistics and supply chain management* (5th ed.). Pearson.

เอกสารอ้างอิง (ต่อ)

- El Baz, J., & Ruel, S. (2021). Can supply chain risk management practices mitigate the disruption impacts on supply chain resilience and robustness? Evidence from an empirical survey in a COVID-19 outbreak era. *International Journal of Production Economics*, 233, 107972.
- Helmold, M., Küçük Yılmaz, A., Dathe, T., & Flouris, T. G. (2022). *Supply chain risk management: Cases and industry insights*. Springer.
- Ho, W. (2008). Integrated analytic hierarchy process and its applications – A literature review. *European Journal of Operational Research*, 186(1), 211–228.
<https://doi.org/10.1016/j.ejor.2007.01.004>
- Hou, Y., Khokhar, M., Zia, S., & Sharma, A. (2022). Assessing the best supplier selection criteria in supply chain management during the COVID-19 pandemic. *Frontiers in Psychology*, 12, 804954. <https://doi.org/10.3389/fpsyg.2021.804954>.
- Humphreys, P., Wong, Y. K., & Chan, F. T. S. (2003). Integrating environmental criteria into the supplier selection process. *Journal of Materials Processing Technology*, 138(1–3), 349–356. [https://doi.org/10.1016/S0924-0136\(03\)00097-9](https://doi.org/10.1016/S0924-0136(03)00097-9)
- Krause, D. R., Handfield, R. B., & Scannell, T. V. (1998). An empirical investigation of supplier development: Reactive and strategic processes. *Journal of Operations Management*, 17(1), 39–58. [https://doi.org/10.1016/S0272-6963\(98\)00030-8](https://doi.org/10.1016/S0272-6963(98)00030-8)
- Kundu, O., James, A. D., & Rigby, J. (2020). Public procurement and innovation: A systematic literature review. *Science and Public Policy*, 47(4), 490-502.
<https://doi.org/10.1093/scipol/scaa029>.
- Leenders, M. R., Johnson, P. F., Flynn, A. E., & Fearon, H. E. (2006). *Purchasing and supply management* (13th ed.). McGraw-Hill.
- Lysons, K., & Farrington, B. (2020). *Procurement and supply chain management* (10th ed.). Pearson.

เอกสารอ้างอิง (ต่อ)

- Monczka, R. M., Handfield, R. B., Giunipero, L. C., & Patterson, J. L. (2020). *Purchasing and supply chain management* (7th ed.). Cengage Learning.
- Saaty, T. L., & Vargas, L. G. (2012). *Models, methods, concepts & applications of the analytic hierarchy process* (2nd ed.). Springer.
- Saaty, T. L. (2016). The analytic hierarchy and analytic network processes for the measurement of intangible criteria and for decision-making. In J. S. Dyer (Ed.), *Multiple criteria decision analysis: State of the art surveys* (2nd ed., pp. 363-419). Springer. https://doi.org/10.1007/978-1-4939-3094-4_10.
- Slack, N., Brandon-Jones, A., & Burgess, N. (2022). *Operations and process management: Principles and practice for strategic impact* (6th ed.). Pearson.
- van Weele, A. J. (2018). *Purchasing and supply chain management: Analysis, strategy, planning and practice* (7th ed.). Cengage Learning.



ส่วนที่ 3: การวิเคราะห์เชิงพยากรณ์

Part 3: Predictive Analytics

“องค์กรที่คาดการณ์อนาคตได้ดีกว่า
ย่อมตัดสินใจได้ดีกว่า”



บทที่ 5

การจัดการและพยากรณ์ความต้องการสินค้า Demand Management and Forecasting

หัวข้อ

- 5.1 การจัดการความต้องการสินค้า
- 5.2 การพยากรณ์ความต้องการสินค้าและรูปแบบข้อมูลอนุกรมเวลา
- 5.3 การพยากรณ์ด้วยการวิเคราะห์อนุกรมเวลาแบบคลาสสิก
- 5.4 การพยากรณ์ด้วยการวิเคราะห์การถดถอยอัตโนมัติ
- 5.5 การประเมินความคลาดเคลื่อนของการพยากรณ์

การจัดการความต้องการสินค้าหรือการจัดการอุปสงค์เป็นกระบวนการที่จำเป็นในการจัดการซัพพลายเชนด้วยเหตุผลหลักหลายประการเช่น (1) การปรับสมดุลระหว่างอุปสงค์และอุปทาน (balancing supply and demand) เพราะการจัดการความต้องการสินค้าอย่างมีประสิทธิภาพช่วยให้ธุรกิจสามารถปรับการผลิตและการจัดหาวัตถุดิบให้ตรงกับความต้องการของตลาดได้อย่างเหมาะสม ลดปัญหาสินค้าขาดตลาดหรือมีสินค้าคงคลังมากเกินไป ซึ่งทั้งสองสถานการณ์นี้สามารถสร้างความเสียหายทางการเงินได้ (2) การเพิ่มความพึงพอใจของลูกค้า (enhancing customer satisfaction) เมื่อธุรกิจสามารถตอบสนองความต้องการของลูกค้าได้อย่างรวดเร็วและแม่นยำ จะช่วยเพิ่มความพึงพอใจของลูกค้า ส่งผลให้เกิดการซื้อซ้ำและการบอกต่อในเชิงบวก นำไปสู่การรักษาฐานลูกค้าและการเติบโตของธุรกิจในระยะยาว (3) การลดต้นทุน (cost reduction) เพราะการจัดการความต้องการสินค้าที่ดีช่วยให้ธุรกิจสามารถวางแผนการผลิต การจัดซื้อ และการจัดส่งได้อย่างมีประสิทธิภาพ ลดการสูญเสียที่เกิดจากการผลิตมากเกินไปหรือการจัดเก็บสินค้าคงคลังที่ไม่จำเป็น นอกจากนี้ยังช่วยลดต้นทุนที่เกี่ยวข้องกับการบริหารจัดการสต็อกและการขนส่ง (4) การเพิ่มความคล่องตัวและความยืดหยุ่น (increasing agility and flexibility) การจัดการความต้องการที่ดีทำให้ธุรกิจสามารถตอบสนองต่อการเปลี่ยนแปลงของตลาดและพฤติกรรมของลูกค้าได้อย่างรวดเร็ว ไม่ว่าจะเป็นการเปลี่ยนแปลงทางเทคโนโลยี แนวโน้มใหม่ ๆ หรือวิกฤตการณ์ต่าง ๆ ที่อาจเกิดขึ้น (5) การใช้ทรัพยากรอย่างมีประสิทธิภาพ (efficient resource utilization) เมื่อความต้องการสินค้าถูก

คาดการณ์และจัดการได้อย่างถูกต้อง ธุรกิจสามารถวางแผนการใช้ทรัพยากรต่าง ๆ เช่น วัตถุดิบ แรงงาน และเวลา ได้อย่างมีประสิทธิภาพ ลดการสูญเสียและเพิ่มผลผลิต (6) การสร้างความได้เปรียบในการแข่งขัน (competitive advantage) เพราะธุรกิจที่สามารถคาดการณ์และตอบสนองต่อความต้องการของตลาดได้ดีกว่าคู่แข่งย่อมได้เปรียบในการแข่งขัน สามารถครองส่วนแบ่งตลาดและรักษาความเป็นผู้นำในตลาดได้ดีกว่า (7) การตัดสินใจที่มีข้อมูลสนับสนุน (data-driven decision making) เพราะการจัดการความต้องการสินค้าเกี่ยวข้องกับการใช้ข้อมูลในการคาดการณ์และวางแผน การมีข้อมูลที่ถูกต้องและการวิเคราะห์ที่แม่นยำช่วยให้ผู้บริหารสามารถตัดสินใจได้อย่างมั่นใจและมีประสิทธิภาพมากขึ้น โดยในการจัดการความต้องการสินค้าจะมีองค์ประกอบที่สำคัญได้แก่ การพยากรณ์ความต้องการสินค้าซึ่งถือว่าเป็นส่วนหนึ่งและเป็นขั้นตอนสำคัญในกระบวนการจัดการความต้องการสินค้า เพราะข้อมูลที่ได้จากการพยากรณ์จะถูกนำมาใช้ในการวางแผนการผลิต การจัดซื้อ และการบริหารจัดการสินค้าคงคลัง การจัดการความต้องการสินค้าเป็นกระบวนการที่กว้างกว่าและครอบคลุมการพยากรณ์ความต้องการสินค้า การพยากรณ์เป็นเครื่องมือที่ช่วยสนับสนุนการจัดการความต้องการให้มีประสิทธิภาพและสอดคล้องกับเป้าหมายขององค์กร ดังนั้นทั้งสองกระบวนการจึงมีความสำคัญและต้องทำงานร่วมกันเพื่อให้ซัพพลายเชนทำงานได้อย่างมีประสิทธิภาพอย่างยั่งยืน (Rajani et al. , 2022, Zhuo, 2019)

5.1 การจัดการความต้องการสินค้า

ความต้องการสินค้าในบริบทของซัพพลายเชนหมายถึงคำสั่งซื้อจริงที่ลูกค้าสั่งซื้อสินค้า ข้อมูลเกี่ยวกับความต้องการสินค้าเป็นสิ่งที่สำคัญสำหรับการวางแผนและบริหารจัดการซัพพลายเชนอย่างมีประสิทธิภาพ หากข้อมูลเกี่ยวกับความต้องการไม่แม่นยำ บริษัทอาจพบความท้าทายในการวางแผนการผลิตและควบคุม ตัวอย่างเช่น หากข้อมูลความต้องการสินค้าที่ไม่แม่นยำถูกส่งต่อไปยังซัพพลายเชนจากล่างสู่บน อาจเกิดการสูญเสียเปล่าในการวางแผนการผลิตและการเตรียมคำสั่งได้ ซึ่งส่งผลให้เกิดปรากฏการณ์เส้มาที่ไม่พึงประสงค์

การบริหารจัดการความต้องการสินค้าในซัพพลายเชนมุ่งเน้นที่การเพิ่มการมองเห็นในความต้องการ (demand visibility) ความคาดการณ์ และความเชื่อถือได้ (Shapiro, 2021) เพื่อให้บริษัทสามารถออกแบบและส่งมอบผลิตภัณฑ์และบริการที่ตอบสนองความต้องการของลูกค้าได้อย่างมีประสิทธิภาพ การบริหารจัดการความต้องการถือเป็นความสามารถในองค์กรที่เฉพาะเจาะจง ซึ่งประกอบด้วยชุดกลยุทธ์และกระบวนการที่กำหนดไว้สำหรับบริษัทที่ผลิตสินค้าและบริการ

การบริหารจัดการความต้องการสินค้าสามารถแบ่งเป็น 4 ขั้นตอนหลักได้แก่ *sense* (ตระหนักถึง), *predict* (คาดการณ์), *seize* (รับสิ่งที่ได้), และ *stabilize* (ทำให้มั่นคง) หรือโมเดล SPSS ของการบริหารจัดการความต้องการในซัพพลายเชน (Teece, 2007; Waller & Fawcett, 2013)

1. **ขั้นตอน Sense (ตระหนักถึง)** ในขั้นตอนนี้บริษัทตรวจจับแนวโน้มในตลาด ทำความเข้าใจและจับความสำคัญของความต้องการของลูกค้า (Ivanov, 2021) จึงสามารถทำนายการเปลี่ยนแปลงที่จะเกิดขึ้นในอนาคตและตระหนักถึงโอกาสในด้านความต้องการได้ เช่น การวิจัยตลาด การวิเคราะห์รูปแบบอารมณ์จากสื่อสังคม และการวิเคราะห์คู่แข่ง เป็นต้น ตัวอย่างเช่น Amazon ใช้เทคโนโลยี big data และ AI เพื่อรับรู้แนวโน้มความต้องการของลูกค้าโดยการวิเคราะห์ข้อมูลการค้นหา การสั่งซื้อ และการรีวิวสินค้า
2. **ขั้นตอน Predict (คาดการณ์)** หลังจากนั้นบริษัทรวบรวมข้อมูลที่เกี่ยวข้องกับความต้องการที่เป็นไปได้ทั้งหมด เพื่อทำการพยากรณ์ความต้องการสินค้า โดยที่การพยากรณ์บางครั้งอาจไม่แม่นยำ การจำลองสถานการณ์สามารถช่วยให้ผู้ตัดสินใจตัดสินใจวางแผนที่ดีที่สุดหรือเตรียมตัวรับมือกับสถานการณ์ที่แย่ที่สุดได้ ตัวอย่างเช่น Amazon ใช้อัลกอริทึมการพยากรณ์เพื่อคาดการณ์แนวโน้มความต้องการในอนาคต เช่น การเพิ่มขึ้นของความต้องการในช่วงเทศกาลหรือสินค้าที่ได้รับความนิยม
3. **ขั้นตอน Seize (จับ)** ขั้นตอนถัดไปคือการปรับทิศทางและปรับปรุงการกำหนดค่าซัพพลายเชน จัดสรร และทำการตัดสินใจใช้ทรัพยากรเพื่อตอบสนองความต้องการของลูกค้า ซึ่งบริษัทต่างๆ อาจมีกลยุทธ์และวิธีการต่างกันในการตอบสนองความต้องการในมุมมองตามวัตถุประสงค์ธุรกิจของตนเอง เช่น เมื่อลูกค้าต้องการผลิตภัณฑ์ที่สามารถปรับแต่งได้ บริษัทที่ขายผลิตภัณฑ์นั้นๆ จะต้องสามารถปรับทิศทางและปรับปรุงซัพพลายเชนเพื่อให้สามารถตอบสนองอย่างรวดเร็วและมีความยืดหยุ่นเพียงพอ ตัวอย่างเช่น เมื่อคาดการณ์ความต้องการสินค้าของลูกค้าที่เพิ่มขึ้น Amazon จะปรับแผนการจัดหาสินค้าและเพิ่มสต็อกในคลังสินค้าเพื่อรองรับความต้องการที่สูงขึ้นนั้นๆ
4. **ขั้นตอน Stabilize (ทำให้มั่นคง)** ขั้นตอนสุดท้ายของโมเดล SPSS คือการทำให้โครงสร้างและกระบวนการในซัพพลายเชนให้มีความมั่นคงและเป็นมาตรฐาน โดยการพิจารณาว่าบริษัทควรทำอย่างไรเพื่อตอบสนองต่อความต้องการของลูกค้าอย่างประสบความสำเร็จ

Dolgui et al., 2020) บางครั้งบริษัทอาจจะต้องพัฒนาและเสริมสร้างความสามารถเพิ่มเติมเพื่อรับมือกับการเปลี่ยนแปลงตามความต้องการที่เปลี่ยนไป ตัวอย่างเช่น ผู้ผลิตนาฬิกาแบบเก่าอาจต้องปรับทิศทางและพัฒนาความแข็งแกร่งในตลาดสมาร์ทวอตช์ เนื่องจากมีแนวโน้มที่ลูกค้ามีความนิยมให้อุปกรณ์สวมใส่มีความฉลาดและสามารถทำอะไรได้เพิ่มขึ้นกว่าเดิม นอกจากนี้องค์กรอาจใช้กลยุทธ์ต่างๆเพื่อกำหนดรูปแบบของความต้องการของลูกค้าให้เป็นไปตามที่ต้องการ (เช่น การใช้กลยุทธ์ราคา กลยุทธ์สร้างความหลากหลายของผลิตภัณฑ์ หรือ กลยุทธ์ปรับเปลี่ยนเวลาและสถานที่การจัดส่ง) อย่างเหมาะสมกับข้อจำกัดทางด้านการผลิตสอดคล้องกับเป้าหมายด้านประสิทธิภาพของบริษัท ตัวอย่างเช่น Amazon ใช้ระบบการจัดการสินค้าคงคลังและการจัดส่งที่มีประสิทธิภาพเพื่อให้แน่ใจว่าสินค้าพร้อมส่งและลดความเสี่ยงจากการขาดแคลนสินค้า

5.2 การพยากรณ์ความต้องการสินค้าและรูปแบบข้อมูลอนุกรมเวลา

5.2.1 การพยากรณ์ความต้องการสินค้า เป็นกิจกรรมที่สำคัญในการบริหารจัดการความต้องการ เนื่องจากการพยากรณ์เป็นพื้นฐานในการวางแผนการผลิตและควบคุมการบริหารซัพพลายเชนให้มีประสิทธิภาพ (Hyndman & Athanasopoulos, 2021) ถึงแม้ว่าการพยากรณ์ความต้องการสามารถทำได้หลากหลายวิธีการ แต่วัตถุประสงค์สูงสุดของการทำการพยากรณ์มีทิศทางเดียวกัน คือ เพื่อเข้าใจและทำนายคำสั่งซื้อจากลูกค้าในอนาคต การพยากรณ์ความต้องการสามารถใช้วิธีการทั้งเชิงคุณภาพและเชิงปริมาณให้เหมาะสมกับระยะเวลาวางแผนระยะสั้น กลาง และยาว

โดยทั่วไปวิธีการคาดการณ์ความต้องการเชิงคุณภาพจะใช้ประสบการณ์และการวินิจฉัยของผู้เชี่ยวชาญ มากกว่าข้อมูลตัวเลขและแบบจำลองทางสถิติ วิธีการที่ใช้กันอย่างแพร่หลายได้แก่

1. **ประสบการณ์ส่วนตัว** ของผู้จัดการซัพพลายเชนโดยเฉพาะในธุรกิจขนาดเล็กและกลาง (SMEs) ใช้ประสบการณ์และสัญชาตญาณของตนเองในการประมาณความต้องการในอนาคต ตัวอย่างเช่น ผู้จัดการร้านพิมพ์ขนาดเล็กหรือร้านพิซซ่าอาจใช้ประสบการณ์ที่ทำงานมานานในอุตสาหกรรม ประกอบกับความรู้เกี่ยวกับความชอบของลูกค้าในท้องถิ่นเพื่อทำนายความต้องการในอนาคต
2. **ความคิดเห็นของผู้เชี่ยวชาญ** คือ การค้นหาความคิดเห็นและความเข้าใจจากผู้เชี่ยวชาญในวงการที่ปรึกษาหรือบุคคลที่มีความรู้สูงสามารถให้ข้อมูลเชิงคุณภาพที่มี

คุณค่าสำหรับการพยากรณ์ความต้องการ ผู้เชี่ยวชาญสามารถให้มุมมองที่มีความรู้ทางตลาด พฤติกรรมของผู้บริโภคและปัจจัยภายนอกที่มีผลต่อความต้องการ

3. **วิจัยตลาดและการสำรวจ** ได้แก่การใช้เทคนิคการวิจัยตลาดเชิงคุณภาพ เช่น กลุ่มโฟกัส (focus group) การสัมภาษณ์ และการสำรวจ เพื่อรวบรวมข้อมูลเชิงคุณภาพเกี่ยวกับความชอบของลูกค้า ตัวอย่างเช่น แนวโน้มในการซื้อและความต้องการที่รับรู้ วิธีการเหล่านี้ช่วยในการจับข้อมูลคุณภาพที่สามารถใช้ประกอบในการวิเคราะห์ทางเชิงปริมาณ
4. **วิธีเดลไฟ** วิธีนี้เป็นการรวบรวมความคิดเห็นจากผู้เชี่ยวชาญผ่านชุดของแบบสอบถามหรือการสัมภาษณ์ที่มีโครงสร้าง ซึ่งมุ่งหวังที่จะบรรลุความเห็นร่วมสำหรับรูปแบบของความต้องการ
5. **การวิเคราะห์สถานการณ์** คือการพิจารณาสถานการณ์ที่แตกต่างกันและผลกระทบที่อาจเกิดขึ้นต่อความต้องการ โดยการประเมินสถานการณ์ที่แตกต่างกันที่เป็นไปได้ (เช่น การเปลี่ยนแปลงทางเศรษฐกิจ ความรุนแรงของการแข่งขัน) เพื่อให้ธุรกิจสามารถเตรียมการสำหรับผลลัพธ์ที่จะเกิดขึ้นได้

วิธีการเชิงคุณภาพเหล่านี้มีความสำคัญเป็นพิเศษเมื่อข้อมูลย้อนหลังมีอยู่อย่างจำกัดและไม่เพียงพอ หรือเมื่อปัจจัยภายนอกและความเคลื่อนไหวของตลาดมีบทบาทสำคัญในการกำหนดรูปแบบแนวโน้มของความต้องการ วิธีการเหล่านี้ช่วยประกอบกับการใช้วิธีการเชิงปริมาณเพื่อช่วยสนับสนุนการตัดสินใจในการวางแผนความต้องการสินค้าและการบริหารจัดการซัพพลายเชนได้อย่างแม่นยำ

สำหรับวิธีการเชิงปริมาณซึ่งใช้ข้อมูลในการทำนาย ยกตัวอย่างเช่น บันทึกยอดขายในอดีต ข้อมูลทางประชากร รูปแบบฤดูกาลของความต้องการ และสภาพอากาศ เป็นต้น วิธีการเชิงปริมาณเหล่านี้ได้แก่

1. **การวิเคราะห์อนุกรมเวลา** เป็นวิธีการพยากรณ์ที่ใช้ข้อมูลยอดขายย้อนหลังเพื่อทำนายความต้องการสินค้าในอนาคต โดยมีความเชื่อที่ว่าความต้องการในอนาคตมักจะมีแนวโน้มหรือมีรูปแบบเดียวกับยอดขายในอดีต ซึ่งมักจะเป็นจริงเมื่อไม่มีความผันผวนหรือความไม่แน่นอนของตลาดเข้ามาเกี่ยวข้อง
2. **วิธีความสัมพันธ์เชิงสาเหตุ** เป็นวิธีการที่ระบุคุณลักษณะที่สำคัญที่อาจมีผลต่อความต้องการสินค้าในอนาคต เช่น เงื่อนไขอากาศ ตัวอย่างเช่น ในช่วงฤดูร้อน การเพิ่มอุณหภูมิอาจทำให้มีความต้องการสำหรับไอศกรีมอย่างมีนัยสำคัญ นอกจากนี้ยังรวมถึง

การแปรผันทางฤดูกาลและวันหยุดพักผ่อนและโปรโมชั่นซึ่งอาจถูกใช้เป็นคุณลักษณะที่สำคัญในการพยากรณ์ความต้องการของลูกค้าในโมเดลการพยากรณ์เหล่านี้

3. **การเรียนรู้ของเครื่อง** เป็นการใช้วิธีการขั้นสูงในการทำนายและพยากรณ์ความต้องการของลูกค้า เช่น Support Vector Machines และ Random Forests วิธีการทางด้านการวิเคราะห์นี้มักจะให้ผลลัพธ์ที่ดีกว่าวิธีการแบบเก่า (Haselbeck et al., 2022; Makridakis et al., 2022)
4. **การจำลอง** เป็นวิธีการที่ใช้สำหรับพัฒนาการพยากรณ์จากทั้งวิธีการเชิงคุณภาพและปริมาณ เพื่อจำลองผลกระทบจากสภาพที่เป็นไปได้ต่างๆ (what-if scenarios) ช่วยให้ผู้จัดการซัพพลายเชนเตรียมพร้อมกับการจัดการความไม่แน่นอนในซัพพลายและความต้องการสินค้า เช่น การจัดกิจกรรมโปรโมชั่นเพื่อเพิ่มยอดขายหลังจากมีความต้องการลดลงในช่วงเศรษฐกิจซบเซา

แม้ว่าจะมีวิธีการเชิงคุณภาพและเชิงปริมาณมีหลากหลายเทคนิคให้เลือกใช้ ผู้จัดการควรเลือกวิธีการที่เหมาะสมที่สุดตรงกับรูปแบบของความต้องการเพื่อให้การพยากรณ์ที่แม่นยำที่สุด (Gilliland, 2010) อย่างไรก็ตาม จะต้องคำนึงถึงเสมอว่าการพยากรณ์มักจะไม่ได้ถูกต้อง 100% โดยทั่วไปการพยากรณ์ความต้องการสินค้าในระยะสั้นมักจะมีความแม่นยำมากกว่าการพยากรณ์ในระยะยาว และการพยากรณ์มักจะแม่นยำมากขึ้นสำหรับข้อมูลรวม (aggregate data) เมื่อเทียบข้อมูลสินค้ารายชิ้น

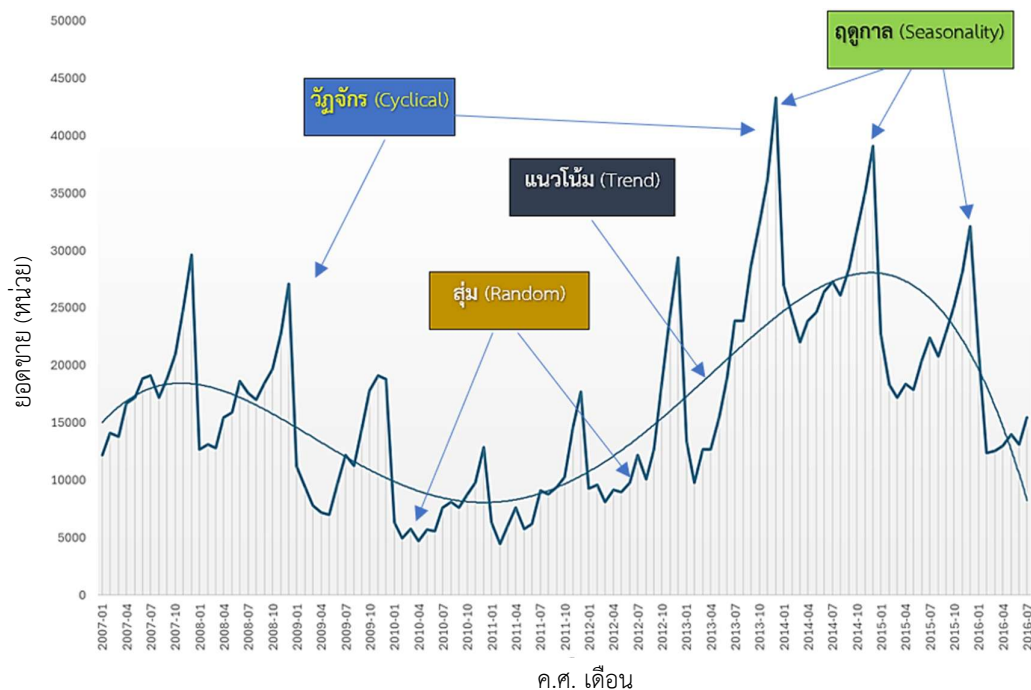
5.2.2 รูปแบบข้อมูลอนุกรมเวลา

โดยทั่วไปแล้ว ข้อมูลอนุกรมเวลาจะถูกบันทึกในช่วงเวลาปกติ วิธีการพยากรณ์อนุกรมเวลามักใช้เวลาเป็นตัวแปรอิสระ และความต้องการที่พยากรณ์โดยวิธีเหล่านี้จะเป็นค่าที่คาดการณ์ในเวลาที่จะเฉพาะเจาะจง ข้อมูลอนุกรมเวลาอาจประกอบด้วยองค์ประกอบหลักสี่ประการ (ดูภาพที่ 5.1) ซึ่งสามารถสำรวจและระบุได้โดยโมเดลการพยากรณ์ที่เหมาะสม

1. **แนวโน้ม (Trend)** คือ แนวโน้มขึ้นหรือลงในอนุกรมเวลา เช่น การเติบโตหรือการลดลงของยอดขายที่ผ่านมา

2. ความเป็นฤดูกาล (Seasonality) เป็นรูปแบบระยะสั้นที่เกิดขึ้นในช่วงเวลาที่เฉพาะเจาะจงและเกิดซ้ำไม่สิ้นสุด เช่น การเพิ่มขึ้นของยอดขายไอศกรีมในช่วงฤดูร้อน อย่างไรก็ตาม ความเป็นฤดูกาลไม่จำเป็นต้องเกี่ยวข้องกับฤดูกาลทางธรรมชาติ สำหรับตัวอย่าง ยอดขายเครื่องดื่มแอลกอฮอล์อาจเพิ่มขึ้นทุกวันหยุดสุดสัปดาห์ ซึ่งก็เป็นรูปแบบระยะสั้นที่เกิดขึ้นซ้ำเช่นกัน

- องค์ประกอบเชิงวัฏจักร (Cyclical Component) เป็นการเคลื่อนไหวระยะยาวในข้อมูลอนุกรมเวลาที่ใช้เวลาหลายปีหรือหลายทศวรรษในการปรากฏ เกิดขึ้นเนื่องจากสถานะทางเศรษฐกิจภายนอก เช่น ภาวะถดถอยและการฟื้นตัว
- สัญญาณรบกวน (Noise, Random) คือความแปรผันแบบสุ่มในอนุกรมเวลาที่เกิดจากปัจจัยที่ไม่สามารถทำนายได้และไม่คาดคิด ซึ่งเกิดขึ้นอย่างไม่สม่ำเสมอ สัญญาณรบกวนมักเป็นสาเหตุของรูปร่างซิกแซกในแผนภาพอนุกรมเวลา



ภาพที่ 5.1 องค์ประกอบของอนุกรมเวลา

ที่มา: ดัดแปลงจาก (Liu, 2022)

5.3 การพยากรณ์ด้วยการวิเคราะห์อนุกรมเวลาแบบคลาสสิก

ดังที่ได้กล่าวไว้ก่อนหน้านี้ การวิเคราะห์อนุกรมเวลาใช้ข้อมูลในอดีตเพื่อการพยากรณ์ การวิเคราะห์ดังกล่าวประกอบด้วยวิธีการต่างๆ ที่สามารถดึงรูปแบบและลักษณะจากข้อมูลอนุกรมเวลาในอดีตมาใช้ในการทำนายความต้องการในอนาคต ในส่วนนี้เราจะพิจารณาวิธีการพยากรณ์อนุกรมเวลาที่คลาสสิกและได้รับความนิยมมากที่สุดบางวิธี ได้แก่

5.3.1 วิธีถัวเฉลี่ยเคลื่อนที่ (Moving Average)

วิธีถัวเฉลี่ยเคลื่อนที่ (Moving Average - MA) สร้างชุดของค่าเฉลี่ยในอนุกรมเวลาภายในช่วงเวลาที่กำหนด (Montgomery et al., 2015) โดยสมมติว่าการสังเกตที่เกิดขึ้นใกล้เคียงกันในเวลา มีค่าใกล้เคียงกัน ดังนั้นการพยากรณ์สำหรับช่วงเวลาถัดไปจะเท่ากับค่าเฉลี่ยของช่วงเวลา N ล่าสุด (ดูสมการ 5.1 ด้านล่าง) การใช้ค่าเฉลี่ยสามารถตัดสัญญาณรบกวนรบกวนออกไปได้

$$F_{t+1} = \frac{A_t + A_{t-1} + \dots + A_{t-N+1}}{N} \quad (5.1)$$

โดยที่ F_{t+1} หมายถึงการพยากรณ์สำหรับช่วงเวลาถัดไป, A_t คือความต้องการจริงล่าสุด, และ N คือช่วงเวลาที่กำหนด

เป็นข้อสังเกตว่าถ้า N ที่มีค่าน้อยจะทำให้การพยากรณ์ตอบสนองได้รวดเร็วกว่า (เช่น ใกล้เคียงกับความต้องการจริงล่าสุด) ในขณะที่ N ที่มีค่ามากจะทำให้การพยากรณ์มีความเสถียรมากขึ้น โดยที่การกำหนดช่วงเวลา N สามารถพิจารณาจากการทดสอบความคลาดเคลื่อน เช่น ค่า MAD และ ค่า MSE

ตัวอย่าง 5.1 การพยากรณ์ด้วยวิธี Moving Average

สมมติว่าเรามีข้อมูลยอดขายของสินค้าหนึ่งในช่วง 5 เดือนที่ผ่านมา ดังนี้:

มกราคม: 100 หน่วย
กุมภาพันธ์: 120 หน่วย
มีนาคม: 130 หน่วย
เมษายน: 150 หน่วย
พฤษภาคม: 170 หน่วย

ถ้าเราต้องการพยากรณ์ยอดขายในเดือนมิถุนายนโดยใช้การพยากรณ์แบบ Moving Average 3

เดือน ของเดือนมิถุนายน (MA 3) เราสามารถคำนวณได้ดังต่อไปนี้

คำนวณค่าเฉลี่ยเคลื่อนที่ของ 3 เดือนล่าสุด:

ค่าพยากรณ์ของเดือนมิถุนายน $= (130 + 150 + 170) / 3 = 450 / 3 = 150$ หน่วย

ดังนั้น การพยากรณ์ยอดขายสำหรับเดือนมิถุนายนคือ 150 หน่วย

5.3.2 วิธีปรับเรียบเอกซ์โปเนนเชียล (Exponential Smoothing)

อีกวิธีการพยากรณ์คลาสสิกที่เรียกว่า วิธีการปรับเรียบเอกซ์โปเนนเชียล (Exponential Smoothing, ES) แทนที่จะใช้ค่าน้ำหนักเท่ากันแบบเดียวกับวิธี MA วิธี ES จะกำหนดค่าน้ำหนักที่ลดลงแบบทวีคูณตามช่วงเวลา (ดูสมการ 5.2 ด้านล่าง)

$$F_{t+1} = \alpha A_t + (1 - \alpha) F_t \quad (5.2)$$

โดยที่ F_t หมายถึงการพยากรณ์ครั้งก่อนหน้า และ $F_t = A_t$ เมื่อ $t=0$ α คือค่าคงที่ในการปรับเรียบ และ $0 \leq \alpha \leq 1$ การใช้ค่าคงที่ในการปรับเรียบ α สามารถช่วยลดสัญญาณรบกวนที่มีความถี่สูงในข้อมูล ค่า α ที่น้อยกว่าจะทำให้การพยากรณ์มีเสถียรภาพมากขึ้น ในขณะที่ค่า α ที่มากกว่าจะทำให้การพยากรณ์ตอบสนองต่อการเปลี่ยนแปลงได้มากขึ้น (Hyndman et al., 2008)

วิธีการพยากรณ์แบบการปรับเรียบเอกซ์โปเนนเชียล (ES) มีข้อดีคือสามารถทำการพยากรณ์ในช่วงเวลาถัดไปได้โดยใช้ข้อมูลในอดีตเพียงเล็กน้อย อย่างไรก็ตาม ควรทราบว่าวิธีนี้จะสามารถพยากรณ์ได้เพียงหนึ่งช่วงเวลข้างหน้าเท่านั้น ทำให้การพยากรณ์ระยะยาวจะไม่มี ความเปลี่ยนแปลง

ตัวอย่าง 5.2 การพยากรณ์ด้วยวิธี Exponential Smoothing

สมมติว่าเรามีข้อมูลยอดขายในช่วง 10 เดือนที่ผ่านมา (ในหน่วย):

มกราคม: 200

กุมภาพันธ์: 215

มีนาคม: 210

เมษายน: 225

พฤษภาคม: 220

มิถุนายน: 230

กรกฎาคม: 240

สิงหาคม: 250

กันยายน: 260

ตุลาคม: 255

ถ้าเราต้องการพยากรณ์ยอดขายในเดือนพฤศจิกายนโดยใช้ Exponential Smoothing โดยกำหนดค่า α (น้ำหนักของการปรับเรียบ) เท่ากับ 0.3 และยอดการพยากรณ์เริ่มต้น (สำหรับเดือนมกราคม) เท่ากับ 200 หน่วย

ขั้นตอนการคำนวณ

1. เริ่มต้นการพยากรณ์: กำหนดยอดการพยากรณ์สำหรับเดือนมกราคมเป็น 200 หน่วย (สมมติให้เท่ากับยอดขายจริง)

2. คำนวณยอดการพยากรณ์สำหรับเดือนกุมภาพันธ์:

$$F_2 = \alpha \times A_1 + (1 - \alpha) \times F_1$$

$$F_2 = 0.3 \times 200 + 0.7 \times 200 = 200 \text{ หน่วย}$$

3. คำนวณยอดการพยากรณ์สำหรับเดือนมีนาคม:

$$F_3 = \alpha \times A_2 + (1 - \alpha) \times F_2$$

$$F_3 = 0.3 \times 215 + 0.7 \times 200 = 204.5 \text{ หน่วย}$$

4. คำนวณยอดการพยากรณ์สำหรับเดือนเมษายนถึงตุลาคมโดยใช้วิธีเดียวกัน:

ตัวอย่างสำหรับเดือนเมษายน:

$$F_4 = 0.3 \times 210 + 0.7 \times 204.5 = 206.4 \text{ หน่วย}$$

5. คำนวณยอดการพยากรณ์สำหรับเดือนพฤศจิกายน:

ทำการคำนวณต่อเนื่องสำหรับแต่ละเดือนจนถึงเดือนตุลาคม จากนั้นใช้ค่า α และยอดขายจริงของเดือนตุลาคมในการคำนวณยอดพยากรณ์สำหรับเดือนพฤศจิกายน

$$F_{11} = \alpha \times A_{10} + (1 - \alpha) \times F_{10} = 0.3 \times 255 + 0.7 \times 240.8 = 245.1 \text{ หน่วย}$$

5.4 การวิเคราะห์อนุกรมเวลาด้วยโมเดลวิเคราะห์การถดถอยแบบอัตโนมัติ

(Autoregressive Model -AR)

โมเดล AR มีความคล้ายคลึงกับโมเดลวิเคราะห์การถดถอย โดยจะเรียนรู้รูปแบบพฤติกรรมของข้อมูลในอดีตเพื่อพยากรณ์แนวโน้มในอนาคต ในโมเดล AR การพยากรณ์ในช่วงเวลาถัดไปที่ t สามารถได้รับผลกระทบจากการสังเกตการณ์ในช่วงเวลา p ก่อนหน้า ซึ่งสามารถแสดงดังสมการ (5.3):

$$y_t = c + \phi_1 y_{t-1} + \phi_2 y_{t-2} + \dots + \phi_p y_{t-p} + \varepsilon_t \quad (5.3)$$

โดยที่ p คืออันดับของโมเดล (เช่น ช่วงเวลาก่อนหน้า) c คือค่าคงที่ ϕ คือพารามิเตอร์ และตัวแปรสุ่ม ε_t คือสัญญาณรบกวนสีขาว (white noise) ตัวอย่างเช่น ในตลาดหุ้น ราคาหุ้นปัจจุบันของบริษัทใดๆ สามารถขึ้นอยู่กับราคาหุ้นก่อนหน้าทั้งหมดในอนุกรมเวลา โมเดล AR คำนวณการถดถอยของค่าที่ผ่านมาของอนุกรมเวลาและพยากรณ์ค่าที่จะเกิดขึ้นในอนาคต (Hyndman & Athanasopoulos, 2018)

5.4.1 โมเดลเคลื่อนที่เฉลี่ย (Moving Average - MA)

โมเดล MA คือกระบวนการ MA ของลำดับ q หรือ $MA(q)$ สามารถได้รับผลกระทบจากปัจจัยภายนอกที่ไม่คาดคิดในช่วงเวลา q ก่อนหน้า (เช่น $t - 1, t - 2, \dots, t - q$) ผลกระทบที่ไม่คาดคิดเหล่านี้มักเรียกกันว่า ข้อผิดพลาดหรือเศษเหลือ ดังนั้น โมเดล MA สามารถแสดงทางคณิตศาสตร์ดังสมการที่ (5.4):

$$y_t = c + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} + \dots + \theta_q \varepsilon_{t-q} + \varepsilon_t \quad (5.4)$$

$$= \mu + \sum_{i=1}^q \theta_i \varepsilon_{t-i} + \varepsilon_t$$

โดยที่ q คือ ลำดับ (เช่น ช่วงเวลาก่อนหน้า) μ คือ ค่าเฉลี่ยของชุดข้อมูล (มักสมมติว่าเท่ากับ 0) θ คือ พารามิเตอร์ และตัวแปรสุ่ม ε_t คือ สัญญาณรบกวนสีขาว

โมเดล ARMA(p, q) จะพิจารณาทั้งสัญญาณรบกวนสุ่ม (ส่วน MA) และส่วนการถดถอยอัตโนมัติ (ส่วน AR) ซึ่งสามารถแสดงดังสมการ (5.5):

$$y_t = c + \sum_{i=1}^p \varphi_i y_{t-i} + \sum_{i=1}^q \theta_i \varepsilon_{t-i} + \varepsilon_t \quad (5.5)$$

5.4.2 โมเดลถดถอยเคลื่อนที่เฉลี่ยอัตโนมัติแบบบูรณาการ (Autoregressive Integrated Moving Average - ARIMA)

อีกหนึ่งโมเดลการพยากรณ์อนุกรมเวลาที่มีความคล้ายคลึงและเป็นที่ยอมรับกันอย่างแพร่หลายคือ โมเดลถดถอยเคลื่อนที่เฉลี่ยอัตโนมัติแบบบูรณาการ (Autoregressive Integrated Moving Average - ARIMA) (Box et al., 2015) โมเดล ARIMA มีประโยชน์ในการสร้างแบบจำลองสำหรับอนุกรมเวลาที่ไม่ใช่รูปแบบฤดูกาล ซึ่งแสดงรูปแบบที่ชัดเจนและไม่ใช้สัญญาณรบกวนสีขาวแบบสุ่ม แต่ถ้าหากมีรูปแบบฤดูกาล เราจำเป็นต้องเพิ่มเงื่อนไขฤดูกาลเข้าไปใน ARIMA หรือที่เรียกว่า SARIMA (Nokeri, 2021)

โมเดล ARIMA(p, d, q) ประกอบด้วยพารามิเตอร์สามตัว ดังนี้:

- p คืออันดับของพจน์ AR (Autoregressive) หรือ จำนวนล่าช้า (lag) ในส่วน AR
- q คืออันดับของพจน์ MA (Moving Average)
- d คือจำนวนครั้งที่ต้องทำการหาความต่าง (differencing)

โมเดล ARIMA(p, d, q) จะพิจารณาทั้งสัญญาณรบกวนสุ่ม (ส่วน MA) และส่วนการถดถอยอัตโนมัติ (ส่วน AR) ซึ่งสามารถแสดงดังสมการ (5.6):

ถ้า $d = 0$: $y_t = Y_t$

ถ้า $d = 1$: $y_t = Y_t - Y_{t-1}$

ถ้า $d = 2$: $y_t = (Y_t - Y_{t-1}) - (Y_{t-1} - Y_{t-2}) = Y_t - 2Y_{t-1} + Y_{t-2}$

$$y_t = c + \varphi_1 y_{t-1} + \varphi_2 y_{t-2} + \cdots + \varphi_p y_{t-p} + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} + \cdots + \theta_q \varepsilon_{t-q} + \varepsilon_t \quad (5.6)$$

โดยที่

Y_t คือ ค่าจริงในช่วงเวลา t

y_t คือ ค่าที่พยากรณ์ในช่วงเวลา t

c คือ ค่าคงที่

φ_t คือ พารามิเตอร์ของส่วน AR

θ_t คือ พารามิเตอร์ของส่วน MA

ε_t คือ white noise

ตัวอย่าง 5.3 การพยากรณ์ด้วยวิธี ARIMA

การคำนวณการพยากรณ์ด้วยวิธี ARIMA ด้วยมือเป็นกระบวนการที่ซับซ้อนและต้องใช้ความเข้าใจเชิงลึกในสถิติและการวิเคราะห์อนุกรมเวลา เนื้อหาด้านล่างนี้เป็นตัวอย่างคร่าว ๆ ของการคำนวณโดยใช้ ARIMA(1,0,1) ซึ่งเป็นหนึ่งในรูปแบบอย่างง่ายรูปแบบหนึ่ง:

ขั้นตอนที่ 1: การกำหนดโมเดล ARIMA

สำหรับ ARIMA (1, 0, 1) ค่าสมการจะเป็น

$$y_t = c + \varphi_1 y_{t-1} + \theta_1 \varepsilon_{t-1} + \varepsilon_t$$

ขั้นตอนที่ 2: การหา Differencing (d ไม่เท่ากับ 0)

ก่อนอื่นทำ differencing ข้อมูล y_t คือ $Y_t - Y_{t-1}$ ถ้าข้อมูลไม่เสถียร

ขั้นตอนที่ 3: การหาประมาณค่า φ และ θ โดยใช้เทคนิคต่างๆเช่น Autocorrelation Function (ACF) และ Partial Autocorrelation Function (PACF)

ขั้นตอนที่ 4: การคำนวณ

สมมติว่ามีข้อมูลต่อไปนี้:

$$y_1=50$$

$$y_2=53$$

$$y_3=55$$

$$\varphi_1 = 0.5, \theta_1 = 0.4, c = 2$$

$$y_4 = 2 + 0.5(55) + 0.4(\varepsilon_3) + \varepsilon_4$$

ขั้นตอนที่ 5: การพยากรณ์ค่าในอนาคต

หลังจากการคำนวณเบื้องต้น คุณสามารถใช้สมการเหล่านี้เพื่อพยากรณ์ค่าในอนาคต โดยแทนค่า $y_{t-1}, \varepsilon_{t-1}$ และค่าที่ได้จากการคำนวณก่อนหน้านี้

5.5 การประเมินความคลาดเคลื่อนของการพยากรณ์

เพื่อประเมินประสิทธิภาพของโมเดลการพยากรณ์ เราสามารถดูที่ความคลาดเคลื่อนของการพยากรณ์ (forecast error) ความคลาดเคลื่อนของการพยากรณ์สามารถแสดงได้ดังนี้: $E_t = A_t - F_t$ โดยที่ E คือ ความคลาดเคลื่อนของการพยากรณ์ A คือ ความต้องการจริง และ F คือ ความต้องการที่พยากรณ์ไว้สำหรับช่วงเวลา t

มีตัวชี้วัดความแม่นยำที่หลากหลายสำหรับการประเมินการพยากรณ์แบบอนุกรมเวลา ซึ่งตัวชี้วัดที่ใช้กันอย่างแพร่หลายได้แก่:

1. ค่าความคลาดเคลื่อนเฉลี่ยกำลังสอง (Mean Squared Error-MSE)

$$MSE = \frac{1}{n} \sum_{t=1}^n E_t^2 \quad (5.7)$$

MSE มีประโยชน์เมื่อความคลาดเคลื่อนของการพยากรณ์มีการกระจายแบบสมมาตรรอบค่าเฉลี่ยศูนย์

2. ค่ารากที่สองของความคลาดเคลื่อนกำลังสองเฉลี่ย (Root Mean Square Error-RMSE)

$$RMSE = \sqrt{\frac{1}{n} \sum_{t=1}^n E_t^2} \quad (5.8)$$

RMSE คือรากที่สองของ MSE โดยที่ RMSE มีความไวต่อค่าผิดปกติมากๆ (outliers) และมีประโยชน์สำหรับการเปรียบเทียบระหว่างโมเดลต่างๆ บนชุดข้อมูลอนุกรมเวลาเดียวกัน

3. ค่าเบี่ยงเบนเฉลี่ยสัมบูรณ์ (Mean Absolute Deviation- MAD)

$$MAD = \frac{1}{n} \sum_{t=1}^n |E_t| \quad (5.9)$$

MAD หรือที่รู้จักกันในชื่อ Mean Absolute Error (MAE) จะรวมค่าความคลาดเคลื่อนในค่าสัมบูรณ์ ดังนั้นความแตกต่างด้านลบจะไม่หักล้างกับความแตกต่างด้านบวก MAD เป็นเครื่องมือที่ดีสำหรับการวัดข้อผิดพลาดที่แท้จริงในการพยากรณ์ อย่างไรก็ตาม เนื่องจากวิธีนี้ให้ค่าเฉลี่ยของข้อผิดพลาดจึงไม่ค่อยมีประโยชน์สำหรับการเปรียบเทียบระหว่างอนุกรมเวลาที่มีสเกลแตกต่างกันอย่างมาก

4. เปอร์เซนต์ค่าเฉลี่ยสัมบูรณ์ (Mean Absolute Percentage Error- MAPE)

$$MAPE = \frac{1}{n} \sum_{t=1}^n \frac{|E_t|}{A_t} \times 100 \quad (5.10)$$

MAPE คล้ายกับ MAD แต่จะพิจารณาความคลาดเคลื่อนของการพยากรณ์สัมพันธ์กับความต้องการจริง และจึงมีประโยชน์สำหรับการเปรียบเทียบระหว่างสองอนุกรมเวลาที่แตกต่างกัน MAPE ยังเป็นเครื่องมือที่ดีเมื่อมีฤดูกาลในอนุกรมเวลาหรือความต้องการเปลี่ยนแปลงอย่างมีนัยสำคัญจากช่วงเวลาหนึ่งไปยังอีกช่วงเวลา MAPE อาจเป็นเมตริกที่นิยมใช้มากที่สุดในการประเมินการพยากรณ์ (Hewamalage et. al., 2023)

ตัวอย่าง 5.5 การประเมินความเที่ยงตรงของการพยากรณ์ด้วย MAD. และ MAPE.

สมมติว่ามีค่าจริงและค่าพยากรณ์ดังนี้:

เดือน	ค่าจริง (Actual Demand)	ค่าพยากรณ์ (Forecast)
1	100	90
2	120	130
3	130	125
4	140	135
5	150	145

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนการคำนวณ MAD:

- หาผลต่างสัมบูรณ์ระหว่างค่าจริงและค่าพยากรณ์แสดงดังตารางต่อไปนี้:

เดือน	ผลต่าง
1	$ 100 - 90 = 10$
2	$ 120 - 130 = 10$
3	$ 130 - 125 = 5$
4	$ 140 - 135 = 5$
5	$ 150 - 145 = 5$

ที่มา: จัดทำโดยผู้แต่ง

- หาค่าเฉลี่ยของผลต่างสัมบูรณ์:

$$MAD = \frac{(10 + 10 + 5 + 5 + 5)}{5} = \frac{35}{5} = 7$$

ดังนั้น MAD ของการพยากรณ์นี้คือ 7

ขั้นตอนการคำนวณ MAPE:

ใช้ข้อมูลเดียวกับตัวอย่าง MAD

1. หาผลต่างสัมบูรณ์ระหว่างค่าจริงและค่าพยากรณ์เป็น % แสดงดังตารางต่อไปนี้:

เดือน	ผลต่าง
1	$\left \frac{100 - 90}{100} \right \times 100 = 10\%$
2	$\left \frac{120 - 130}{120} \right \times 100 = 8.33\%$
3	$\left \frac{130 - 125}{130} \right \times 100 = 3.85\%$
4	$\left \frac{140 - 135}{140} \right \times 100 = 3.57\%$
5	$\left \frac{150 - 145}{150} \right \times 100 = 3.33\%$

ที่มา: จัดทำโดยผู้แต่ง

2. หาค่าเฉลี่ยของ % ทั้งหมด:

$$MAPE = \frac{(10 + 8.33 + 3.85 + 3.57 + 3.33)}{5} = \frac{29.08}{5} = 5.82\%$$

ดังนั้น MAPE ของการพยากรณ์นี้คือ 5.82%

จากตัวอย่างจะเห็นได้ว่า MAD จะวัดความผิดพลาดเฉลี่ยในหน่วยเดียวกับข้อมูล แต่ MAPE จะวัดความผิดพลาดเฉลี่ยเป็นเปอร์เซ็นต์ของค่าจริง ทำให้สามารถเปรียบเทียบกับการพยากรณ์อื่น ๆ ได้ง่าย

สรุป

การจัดการความต้องการสินค้าเป็นกระบวนการที่สำคัญในการจัดการซัพพลายเชนและการดำเนินธุรกิจ โดยมีส่วนช่วยในการปรับสมดุลระหว่างอุปสงค์และอุปทาน การเพิ่มความพึงพอใจของลูกค้า การลดต้นทุน การเพิ่มความคล่องตัวและความยืดหยุ่น การใช้ทรัพยากรอย่างมีประสิทธิภาพ การสร้างความได้เปรียบในการแข่งขัน การสนับสนุนการตัดสินใจที่มีข้อมูลสนับสนุน การจัดการความต้องการสินค้าเกี่ยวข้องโดยตรงกับการใช้ข้อมูลในการคาดการณ์และวางแผน การมีข้อมูลที่ถูกต้องและการวิเคราะห์ที่แม่นยำช่วยให้ผู้บริหารสามารถตัดสินใจได้อย่างมั่นใจและมีประสิทธิภาพมากขึ้น

สำหรับการจัดการความต้องการสินค้า เราสามารถใช้โมเดล SPSS ได้แก่ *sense* (ตระหนักถึง) *predict* (คาดการณ์) *seize* (รับสิ่งที่ได้) และ *stabilize* (ทำให้มั่นคง) การพยากรณ์ความต้องการสินค้าถือว่าเป็นส่วนหนึ่งที่สำคัญของการจัดการความต้องการสินค้า โดยที่สามารถพยากรณ์ได้ 2 ลักษณะได้แก่ การพยากรณ์เชิงคุณภาพและเชิงปริมาณ การพยากรณ์เชิงคุณภาพจะใช้เมื่อไม่มีข้อมูลในอดีต จึงจำเป็นที่จะต้องใช้ความคิดเห็น สัญชาตญาณของผู้เชี่ยวชาญเข้ามาช่วยในการพยากรณ์ได้แก่ การพยากรณ์ความต้องการของสินค้าใหม่ เป็นต้น ส่วนการพยากรณ์เชิงปริมาณใช้เมื่อมีข้อมูลในอดีตที่เพียงพอสำหรับการพยากรณ์ การวิเคราะห์อนุกรมเวลาถือว่าเป็นเครื่องมือสำคัญที่ใช้ในการพยากรณ์เชิงปริมาณ โดยรูปแบบ (pattern) ของข้อมูลอนุกรมเวลาสามารถแบ่งได้เป็น แนวโน้ม ฤดูกาล วัฏจักร หรือแบบผสม สำหรับวิธี การพยากรณ์อนุกรมเวลานั้นสามารถแบ่งได้เป็น วิธีการพยากรณ์แบบคลาสสิกและวิธีสมัยใหม่ วิธีแบบคลาสสิกเป็นวิธีที่ใช้กันมานาน เช่น วิธี Moving Average วิธี Exponential Smoothing วิธี ARIMA ในขณะที่วิธีสมัยใหม่จะประยุกต์ใช้การเรียนรู้ของเครื่อง ได้แก่วิธีที่พัฒนามาจากการใช้อัลกอริทึมโครงข่ายประสาทเทียม (Neural Network) เป็นต้น ซึ่งแบบหลังนั้นเป็นที่นิยมใช้ในปัจจุบันเนื่องจากสามารถพิจารณาปัจจัยร่วมที่มีผลต่อความต้องการสินค้า (โปรโมชัน, อายุ, เพศ, รายได้ และอื่น ๆ) ได้มากกว่าหนึ่งปัจจัยพร้อมกัน แต่อย่างไรก็ตามกลับมีความยุ่งยากซับซ้อนในการคำนวณมากกว่าวิธีแรก (Areerakulkan et. al., 2024) อย่างไรก็ตามสำหรับการพยากรณ์ในทุกๆครั้งจะต้องยึดตามหลักการพยากรณ์ในข้อที่ว่าไม่มีการพยากรณ์ใดที่ถูกต้อง 100% ดังนั้นเมื่อมีการพยากรณ์ในทุกๆกรณีจะต้องมีการวัดค่าความคลาดเคลื่อน (error) ของการพยากรณ์ประกอบกันด้วยเสมอ โดยวิธีในการวัดค่าความผิดพลาดการพยากรณ์จะมีหลายวิธีด้วยกัน ยกตัวอย่างเช่น MSE RMSE MAD และ MAPE แต่ละวิธีล้วนมีจุดแข็งและจุดอ่อนที่แตกต่างกันแล้วแต่วัตถุประสงค์ในการใช้งาน อย่างเช่น MAD จะวัดความผิดพลาดเฉลี่ยในหน่วยเดียวกับ

ข้อมูล แต่ MAPE จะวัดความผิดพลาดเฉลี่ยเป็นเปอร์เซ็นต์ของค่าจริง ทำให้สามารถเปรียบเทียบกับ
การพยากรณ์อื่น ๆ ได้ง่าย โดยสรุปแล้วประสิทธิภาพของการจัดการซัพพลายเชนนั้นขึ้นอยู่กับ
การจัดการและการพยากรณ์ความต้องการสินค้าเนื่องจากการจัดการและการพยากรณ์ที่ดีนั้นจะช่วยสร้าง
รายรับ ลดต้นทุน รวมไปถึงการวางแผนและการตัดสินใจที่ถูกต้องและมีประสิทธิภาพ

แบบฝึกหัด

จงตอบคำถามต่อไปนี้ ส่งงานกลุ่มที่ email: natapat.ar@ssru.ac.th

1. อธิบายความสำคัญของการพยากรณ์ความต้องการสินค้าในบริบทของการบริหารจัดการซัพพลายเชน
2. เปรียบเทียบข้อดีและข้อเสียของวิธีการพยากรณ์เชิงคุณภาพ (Qualitative Forecasting) กับวิธีการพยากรณ์เชิงปริมาณ (Quantitative Forecasting)
3. อธิบายความแตกต่างระหว่างวิธีการพยากรณ์เชิงคุณภาพ (Qualitative Forecasting) และวิธีการพยากรณ์เชิงปริมาณ (Quantitative Forecasting) พร้อมยกตัวอย่างสถานการณ์ที่เหมาะสมสำหรับการใช้วิธีการทั้งสองแบบ
4. ให้นักศึกษาวิเคราะห์ข้อมูลยอดขายสินค้าแสดงดังตารางข้างล่างต่อไปนี้ และระบุว่ามีความโน้ม (Trend) หรือฤดูกาล (Seasonality) หรือไม่ จากนั้นให้ใช้ข้อมูลที่ได้ในการพยากรณ์ยอดขายในปีถัดไป

ปี	ไตรมาส 1	ไตรมาส 2	ไตรมาส 3	ไตรมาส 4
2021	300	350	280	400
2022	320	370	290	420
2023	340	390	300	450

5. บริษัท ABC ต้องการพยากรณ์ยอดขายสินค้าชนิดหนึ่งในไตรมาสหน้า โดยที่ไม่มีข้อมูลยอดขายในอดีตเลย ให้นักศึกษาเลือกวิธีการพยากรณ์ที่เหมาะสมและอธิบายเหตุผลที่เลือกวิธีการนั้น
6. วิเคราะห์ข้อดีและข้อเสียของการใช้วิธีการปรับเรียบเอ็กซ์โปเนนเชียล (Exponential Smoothing) ในการพยากรณ์ยอดขายสินค้า และเปรียบเทียบกับวิธีการพยากรณ์แบบอื่นๆ ที่คุณรู้จัก
7. กรณีศึกษาของบริษัทโทรศัพท์มือถือ มีการออกขายโทรศัพท์มือถือมาแล้ว 6 รุ่น มีข้อมูลยอดขายดังตารางด้านล่าง โดยที่จะมีการปล่อยรุ่นที่ 7 ออกมาในปีหน้า จงทำการพยากรณ์ความต้องการสินค้าของรุ่นที่ 7 (Week 1-12) เพื่อจะนำไปใช้ในการวางแผนการผลิตต่อไป

	Launch #						
	1	2	3	4	5	6	7
Week1	142,000	147,000	162,000	178,000	196,000	235,000	
Week2	125,000	129,000	136,000	142,000	156,000	175,000	
Week3	62,000	69,000	73,000	81,000	98,000	110,000	
Week4	39,000	42,000	47,000	52,000	58,000	62,000	
Week5	11,500	13,000	15,000	19,000	23,000	27,000	
Week6	13,600	14,275	15,250	16,900	17,800	19,000	
Week7	10,950	11,900	12,500	13,750	14,250	15,000	
Week8	6,035	6,180	6,375	6,950	7,600	8,500	
Week9	3,925	4,010	4,190	4,435	4,960	5,500	
Week10	2,795	2,825	2,960	3,045	3,250	3,500	
Week11	1,910	1,975	2,010	2,190	2,255	2,500	
Week12	1,900	1,955	2,000	2,167	2,310	2,500	

8. จากข้อ 7 จงแสดงวิธีในการคำนวณค่าความคลาดเคลื่อนของการพยากรณ์ ดังต่อไปนี้

- MAD
- MSE
- MAPE
- RMSE

และเลือกวิธีที่จะใช้ในการวัดค่าความคลาดเคลื่อน พร้อมทั้งให้เหตุผลประกอบว่าทำไมจึงใช้วิธีนั้น

เอกสารอ้างอิง

- Areerakulkan, N., Moryadee, C., Trakoonsanti, L., Khaengkhan, M., & Setthachotsombut, N. (2024). A new product's demand forecasting using artificial neural network. In *Proceedings of the 26th International Conference on Enterprise Information Systems - Volume 1: ICEIS* (pp. 417-424). SciTePress. <https://doi.org/10.5220/0012734800003690>.
- Box, G. E. P., Jenkins, G. M., Reinsel, G. C., & Ljung, G. M. (2015). *Time series analysis: Forecasting and control* (5th ed.). Wiley.
- Dolgui, A., Ivanov, D. and Sokolov, B. (2020) Reconfigurable Supply Chain: The X-network. *International Journal of Production Research*, 58, 4138-4163. <https://doi.org/10.1080/00207543.2020.1774679>
- Gilliland, M. (2010). The business forecasting deal: Exposing myths, eliminating bad practices, providing practical solutions. Wiley.
- Haselbeck, F., Kuhn, T., & Dimpler, J. (2022). Machine learning outperforms classical forecasting on horticultural sales predictions. *Smart Agricultural Technology*, 2, 100049. <https://doi.org/10.1016/j.atech.2021.100049>.
- Hewamalage, H., Ackermann, K., & Bergmeir, C. (2023). Forecast evaluation for data scientists: Common pitfalls and best practices. *Data Mining and Knowledge Discovery*, 37(2), 788-832. <https://doi.org/10.1007/s10618-022-00894-5>.
- Hyndman, R. J., Koehler, A. B., Ord, J. K., & Snyder, R. D. (2008). *Forecasting with exponential smoothing: The state space approach*. Springer.
- Hyndman, R. J., & Athanasopoulos, G. (2018). *Forecasting: Principles and practice* (2nd ed.). OTexts. <https://otexts.com/fpp2/>.
- Hyndman, R. J., & Athanasopoulos, G. (2021). *Forecasting: Principles and practice* (3rd ed.). OTexts.
- Ivanov, D. (2021). Supply Chain Viability and the COVID-19 Pandemic: A Conceptual and Formal Generalisation of Four Major Adaptation Strategies. *International Journal of Production Research*, 59, 3535-3552. <https://doi.org/10.1080/00207543.2021.1890852>

เอกสารอ้างอิง (ต่อ)

- Liu, K. Y. (2022). *Supply Chain Analytics: Concepts, Techniques and Applications*. (1st ed.) Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-92224-5>.
- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2022). The M5 accuracy competition: Results, findings, and conclusions. *International Journal of Forecasting*, 38(4), 1346–1364. <https://doi.org/10.1016/j.ijforecast.2021.11.013>
- Montgomery, D. C., Jennings, C. L., & Kulahci, M. (2015). *Introduction to time series analysis and forecasting* (2nd ed.). Wiley.
- Nokeri, T. C. (2021). Forecasting using ARIMA, SARIMA, and the additive model. In T. C. Nokeri (Ed.), *Implementing machine learning for finance: A systematic approach to predictive risk and performance analysis for investment portfolios* (pp. 21-50). Apress. https://doi.org/10.1007/978-1-4842-7110-0_2.
- Rajani, R. L., Hegde, G. S., Kumar, R., & Chauhan, P. (2022). Demand management strategies role in sustainability of service industry and impacts performance of company: Using SEM approach. *Journal of Cleaner Production*, 369, 133311. <https://doi.org/10.1016/j.jclepro.2022.133311>.
- Shapiro, J. F. (2021). *Modeling the supply chain* (2nd ed.). Cengage Learning.
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, 28(13), 1319-1350. <https://doi.org/10.1002/smj.640>.
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77-84. <https://doi.org/10.1111/jbl.12010>.
- Zhuo, Z. (2019). Supply Chain from the Demand Orientation: A Systematic Literature Review and Theoretical Model Construction. *Mathematics and Computer Science*, 4(2), 41-56. <https://doi.org/10.11648/j.mcs.20190402.11>

บทที่ 6

การพยากรณ์ความต้องการสินค้า ด้วยวิธีการเรียนรู้ของเครื่อง

Machine Learning Methods for Demand Forecasting

หัวข้อ

- 6.1 แบบจำลองการถดถอยเชิงเส้น (Linear Regression: LR)
- 6.2 แบบจำลองป่าสุ่ม (Random Forest: RF)
- 6.3 โครงข่ายประสาทเทียม (Artificial Neural Network: ANN)

ในช่วงสิบปีที่ผ่านมา ชัฟฟลายเซมมีความซับซ้อนมากขึ้น อันเนื่องมาจากความผันผวนของความต้องการสินค้า การแข่งขันในระดับโลก การเติบโตอย่างรวดเร็วของอีคอมเมิร์ซ และปัญหาที่เกิดจากเหตุการณ์ไม่คาดคิด เช่น การระบาดใหญ่และความขัดแย้งทางภูมิรัฐศาสตร์ ปัจจัยเหล่านี้ทำให้ระบบชัฟฟลายเซมมีความไม่แน่นอนสูงขึ้นและมีความเชื่อมโยงพึ่งพาซึ่งกันและกันมากกว่าเดิม ด้วยเหตุนี้ ผู้จัดการด้านโลจิสติกส์และการจัดการชัฟฟลายเซมจึงไม่สามารถพึ่งพาเพียงประสบการณ์ส่วนบุคคลหรือแบบจำลองทางคณิตศาสตร์พื้นฐาน (เชิงเส้นตรง) ในการตัดสินใจได้อีกต่อไป องค์กรจำเป็นต้องใช้เครื่องมือวิเคราะห์ที่มีประสิทธิภาพสูง ซึ่งสามารถจัดการกับข้อมูลขนาดใหญ่ ความสัมพันธ์ที่ซับซ้อน และรูปแบบข้อมูลที่ซ่อนอยู่ได้ Machine Learning (ML) จึงเป็นหนึ่งในเทคโนโลยีที่สำคัญที่ตอบโจทย์ความต้องการดังกล่าว

ML เป็นสาขาหนึ่งของปัญญาประดิษฐ์ที่มุ่งพัฒนาอัลกอริทึมซึ่งสามารถเรียนรู้จากข้อมูลที่เกิดขึ้นในหน่วยงานจริงในรูปแบบต่างๆ กัน เช่น ข้อมูลยอดขาย เสียง ข้อความ พฤติกรรมผู้ใช้งาน หรือสื่อสังคมออนไลน์ ทำให้การพยากรณ์มีความแม่นยำ รวดเร็วจนอาจถึงการพยากรณ์แบบทันทีทันใด (real-time) (James et al., 2021) โดยหลักการแล้ว ML จะพยายามประมาณฟังก์ชัน $f(x)$ ที่รับตัวแปรนำเข้าหรือปัจจัยนำเข้าได้หลายตัวแปร หลังจากนั้นจะทำการเรียนรู้จากชุดข้อมูลฝึกสอน (training data set) และนำความรู้ที่ได้ไปใช้พยากรณ์ผลลัพธ์ ในบริบทของการพยากรณ์ความต้องการสินค้าในชัฟฟลายเซม ตัวแปรนำเข้าอาจเป็นยอดขายในอดีต ข้อมูลราคา กิจกรรมส่งเสริมการตลาด สภาพอากาศ ระยะเวลา (lead time) หรือตัวชี้วัดเศรษฐกิจมหภาค ส่วนตัวแปรผลลัพธ์อาจเป็นการพยากรณ์ความต้องการสินค้า ระดับสินค้าคงคลัง ความน่าจะเป็นของความเสี่ยงจากผู้ส่งมอบ หรือความล่าช้าในการขนส่ง เป็นต้น (Carbonneau et al., 2008)

ML มีบทบาทสำคัญในซัพพลายเชน และทำงานควบคู่กับ การวิเคราะห์เชิงพยากรณ์ (predictive analytics) ซึ่งมุ่งตอบคำถามว่า “มีแนวโน้มจะเกิดอะไรขึ้นต่อไป” อย่างไรก็ตาม ML ไม่ได้จำกัดอยู่เพียงเพื่อการพยากรณ์เท่านั้น แต่ยังสามารถใช้สนับสนุนการตัดสินใจกำหนดแนวทาง (prescriptive) เพื่อเพิ่มประสิทธิภาพในการตัดสินใจ ตัวอย่างเช่น อาจใช้ อัลกอริทึมหนึ่งใน ML ที่มีชื่อเรียกว่าอัลกอริทึมป่าสุ่ม (Random Forest, RF) เพื่อประมาณยอดขายสินค้า จากนั้นนำยอดขายที่ได้จากการพยากรณ์ไปใช้ในแบบจำลองทางสินค้าคงคลังอย่างเช่น แบบจำลอง Newsvendor เพื่อหาระดับสินค้าคงคลังหรือปริมาณสต็อกสินค้าที่เหมาะสมที่สุดเพื่อลดต้นทุนการจัดการสินค้าคงคลังด้วยวิธีนี้ ML จึงทำหน้าที่ทั้งเป็นเครื่องมือพยากรณ์และเป็นกลไกสนับสนุนการตัดสินใจเพื่อวางกลยุทธ์การจัดการสินค้าคงคลังและปริมาณสินค้าคงคลังที่เหมาะสม

งานวิจัยด้านการจัดการซัพพลายเชนแสดงให้เห็นว่า วิทยาการข้อมูลและการวิเคราะห์ขั้นสูงสามารถเปลี่ยนแปลงทั้งประสิทธิภาพการดำเนินงานและความสามารถในการวางแผนเชิงกลยุทธ์ขององค์กรได้อย่างมีนัยสำคัญ Waller และ Fawcett (2013) ให้เหตุผลว่าการวิเคราะห์เชิงพยากรณ์ (predictive analytics) และ ข้อมูลขนาดใหญ่ (big data) กำลังเปลี่ยนแปลงโครงสร้างและการบริหารซัพพลายเชน ทำให้องค์กรสามารถสร้างความได้เปรียบในการแข่งขันผ่านความสามารถในการประมวลผลข้อมูลที่เหนือกว่า (Areerakulkan & Pongpech, 2021) ในทำนองเดียวกัน พบว่าเทคนิค ML มักให้ผลลัพธ์ดีกว่าวิธีพยากรณ์เชิงสถิติแบบดั้งเดิม โดยเฉพาะในสถานการณ์ที่รูปแบบอุปสงค์มีความไม่เป็นเชิงเส้นและโครงสร้างข้อมูลมีความซับซ้อน ผลการศึกษาดังกล่าวสะท้อนว่า ML ไม่เพียงช่วยเพิ่มความแม่นยำของการพยากรณ์ แต่ยังช่วยยกระดับคุณภาพของการตัดสินใจอีกด้วย

โดยทั่วไป ML สามารถจำแนกได้เป็นสามกลุ่มหลัก กลุ่มแรกคือ การเรียนรู้ภายใต้การกำกับ (supervised learning) ซึ่งอาศัยข้อมูลที่มีตัวแปรผลลัพธ์กำกับ เช่น แบบจำลองการถดถอย ต้นไม้ตัดสินใจ หรือวิธีแบบกลุ่มผสม (ensemble) เช่น Random Forest วิธีเหล่านี้มักใช้ในการพยากรณ์ อุปสงค์ ความเสี่ยง และต้นทุน กลุ่มที่สองคือ การเรียนรู้แบบไม่มีการกำกับ (unsupervised learning) เช่น อัลกอริทึมการจัดกลุ่ม (clustering) ซึ่งมุ่งค้นหารูปแบบที่ซ่อนอยู่ในข้อมูลที่ไม่มีผลลัพธ์กำกับ มักใช้ในการแบ่งกลุ่มลูกค้า การจัดกลุ่มสินค้า หรือการตรวจจับความผิดปกติ กลุ่มที่สามคือ การเรียนรู้เชิงลึก (deep learning) ซึ่งใช้โครงข่ายประสาทเทียมหลายชั้น เหมาะสำหรับการจัดการข้อมูลที่มีมิติสูงและข้อมูลลำดับเวลา เช่น ชุดข้อมูลขนาดใหญ่หรือข้อมูลจากเซนเซอร์แบบเรียลไทม์ในระบบซัพพลายเชนดิจิทัล

แม้ ML จะมีศักยภาพสูง แต่การนำไปใช้ในซัพพลายเชนก็มีความท้าทายหลายประการ คุณภาพข้อมูลยังคงเป็นประเด็นสำคัญ เนื่องจากข้อมูลที่ไม่ถูกต้องหรือไม่ครบถ้วนสามารถทำให้ประสิทธิภาพของแบบจำลองลดลงอย่างมาก ความสามารถในการอธิบายผลลัพธ์ของแบบจำลอง (model interpretability) ก็เป็นอีกประเด็นหนึ่ง (Baryannis et al., 2019) โดยเฉพาะในบริบท

องค์กรที่ผู้ตัดสินใจต้องการเหตุผลรองรับผลลัพธ์จากอัลกอริทึมแต่ผู้ตัดสินใจอาจจะยังไม่มีความรู้หรือความเข้าใจจำกัด นอกจากนี้ การบูรณาการแบบจำลอง ML เข้ากับระบบธุรกิจที่มีอยู่ เช่น ERP หรือแพลตฟอร์มการจัดการขนส่ง ยังต้องอาศัยทั้งความสามารถทางเทคโนโลยีและความพร้อมขององค์กร หากการดำเนินการทางเทคนิคไม่สอดคล้องกับความเข้าใจในเชิงบริหาร องค์กรอาจไม่สามารถได้รับประโยชน์จาก ML อย่างเต็มที่

บทนี้มีวัตถุประสงค์เพื่ออธิบายกรอบแนวคิดและขั้นตอนสำหรับการประยุกต์ใช้ ML ในการพยากรณ์ความต้องการสินค้า โดยเชื่อมโยงหลักการของการเรียนรู้เชิงสถิติและการวิเคราะห์ข้อมูลเข้ากับการคาดการณ์ความต้องการสินค้าในบริบทของการจัดการซัพพลายเชน เนื้อหาครอบคลุมแนวคิดพื้นฐานของแบบจำลอง ML กระบวนการพัฒนาและประเมินแบบจำลอง เพื่อให้ผู้อ่านสามารถประยุกต์ใช้เทคนิคดังกล่าวในการสนับสนุนการวางแผนและเพิ่มประสิทธิภาพของการจัดการซัพพลายเชนได้อย่างมีประสิทธิภาพ

6.1 แบบจำลองการถดถอยเชิงเส้น (Linear Regression: LR)

แบบจำลองการถดถอยเป็นเครื่องมือพื้นฐานและมีประโยชน์ในงานวิเคราะห์เชิงพยากรณ์ และยังเป็นองค์ประกอบสำคัญของทั้งสถิติประยุกต์และ ML ในบริบทของโลจิสติกส์และการจัดการซัพพลายเชน แบบจำลองการถดถอยถูกใช้เพื่อแสดงและพยากรณ์ความสัมพันธ์ระหว่างตัวแปรอิสระกับตัวแปรตาม โดยมีเป้าหมายเพื่อวิเคราะห์ว่าปัจจัยต่าง ๆ ส่งผลกระทบต่อผลลัพธ์ที่สนใจอย่างไร ในซัพพลายเชน ตัวแปรตามอาจเป็นปริมาณความต้องการสินค้า ต้นทุนการขนส่ง ระยะเวลา นำ หรือระดับสินค้าคงคลัง ส่วนตัวแปรอิสระอาจได้แก่ ราคา ยอดขาย ฤดูกาล ระยะทาง ปริมาณคำสั่งซื้อ และตัวชี้วัดทางเศรษฐกิจมหภาค แบบจำลองการถดถอยจึงเป็นเครื่องมือที่เชื่อมโยงเหตุและผล ทำให้องค์กรสามารถตัดสินใจบนพื้นฐานของข้อมูลได้อย่างเป็นระบบ

6.1.1 การถดถอยเชิงเส้นอย่างง่าย

Simple Linear Regression (SLR) หรือ การถดถอยเชิงเส้นอย่างง่าย เป็นแบบจำลองทางสถิติที่ใช้ศึกษาความสัมพันธ์เชิงเส้นระหว่างตัวแปรอิสระ 1 ตัว (X) และตัวแปรตาม 1 ตัว (Y) โดยมีวัตถุประสงค์หลัก 2 ประการ ตัวอย่างเช่น วิเคราะห์ผลกระทบของราคา ต่อ ความต้องการสินค้า หรือ วิเคราะห์ผลกระทบของ ระยะทาง ต่อ ต้นทุนขนส่ง สมการถดถอยเชิงเส้นอย่างง่ายแสดงดังสมการที่ 6.1

$$y = \beta_0 + \beta_1 x + \epsilon \quad (6.1)$$

โดยที่

y = ตัวแปรตาม (Dependent Variable)

x = ตัวแปรอิสระ (Independent Variable)

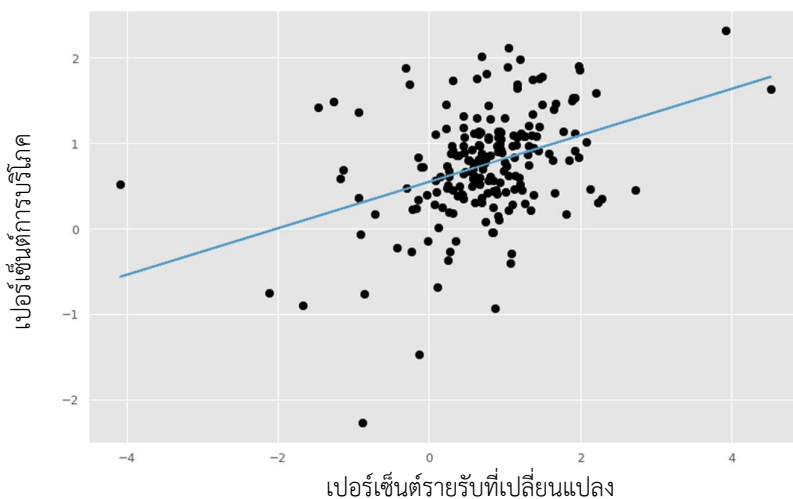
β_0 = ค่าคงที่ (Intercept)

β_1 = ความชัน (Slope)

ϵ = ค่าความคลาดเคลื่อน (Error)

สมมติฐานของการถดถอยเชิงเส้นอย่างง่ายมีหลายประการ เพื่อให้การประมาณค่าและการอนุมานทางสถิติมีความถูกต้อง ได้แก่

สมมติฐานที่ 1 ความเป็นเชิงเส้นตรง (linearity) ซึ่งระบุว่าค่าคาดหวังของ y จะต้องเป็นฟังก์ชันเชิงเส้นตรงของ x โดยประมาณเท่านั้น ในกรณีที่ความสัมพันธ์ไม่สามารถประมาณเป็นเชิงเส้นตรงได้ ไม่ควรใช้สมการถดถอยเชิงเส้นอย่างง่าย (แสดงดังสมการที่ 6.1) ในการประมาณค่า เพราะจะส่งผลให้การประมาณค่าไม่เที่ยงตรงและมีความคลาดเคลื่อนสูง แสดงดังภาพที่ 6.1 จะเห็นได้ว่า เปอร์เซ็นต์รายรับที่เปลี่ยนแปลงจะมีแนวโน้มความสัมพันธ์กับเปอร์เซ็นต์การบริโภคที่เพิ่มขึ้นลดลงโดยประมาณเป็นเส้นตรง



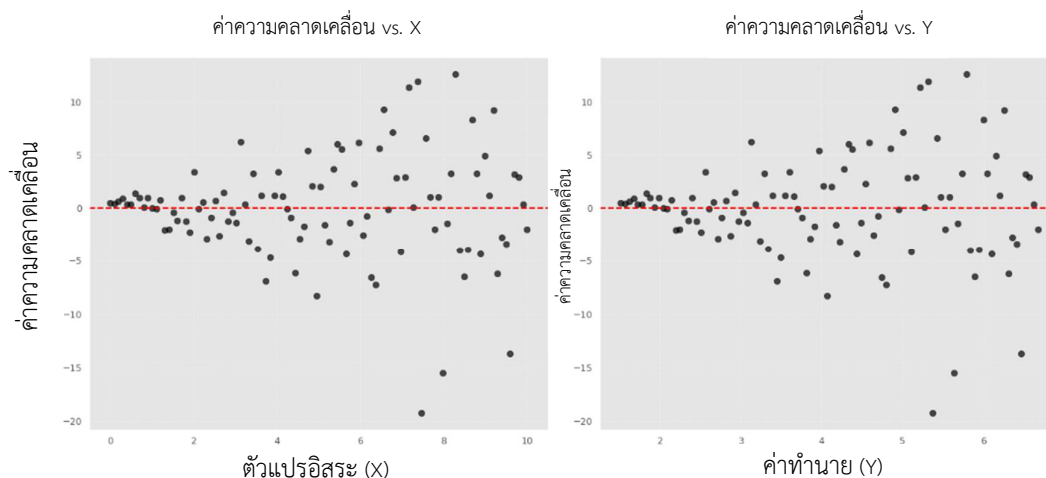
ภาพที่ 6.1 ความสัมพันธ์เชิงเส้นตรง

ที่มา: ภาพโดยผู้แต่ง

สมมติฐานที่ 2 ความเป็นอิสระของค่าคลาดเคลื่อน (independence of errors) ซึ่งกำหนดให้ค่าคลาดเคลื่อนไม่มีความสัมพันธ์กันระหว่างแต่ละตัวอย่าง หมายถึง ค่าคลาดเคลื่อนของข้อมูลแต่ละหน่วยต้องไม่สัมพันธ์กัน หรือความผิดพลาดของการพยากรณ์ในจุดหนึ่งไม่ควรส่งผลต่อความผิดพลาดของอีกจุดหนึ่ง ในเชิงคณิตศาสตร์สามารถเขียนได้ว่า $Cov(\varepsilon_i, \varepsilon_j) = 0$ เมื่อ $i \neq j$ ซึ่งหมายความว่าค่าคลาดเคลื่อนของข้อมูลแต่ละหน่วยเกิดขึ้นอย่างอิสระต่อกัน สมมติฐานนี้มีความสำคัญต่อความน่าเชื่อถือของการวิเคราะห์ทางสถิติ หากค่าคลาดเคลื่อนมีความสัมพันธ์กัน เช่น เกิดค่าบวกหรือลบต่อเนื่องหลายช่วงเวลา (autocorrelation) อาจแสดงว่าแบบจำลองยังไม่สามารถอธิบายโครงสร้างของข้อมูลได้ครบถ้วน และอาจทำให้ผลการทดสอบทางสถิติหรือการพยากรณ์มีความคลาดเคลื่อน

สมมติฐานที่ 3 ค่าคลาดเคลื่อนมีค่าเฉลี่ยเป็นศูนย์ (zero mean of errors) หมายความว่าโดยเฉลี่ยแล้วส่วนคลาดเคลื่อน ε จะไม่เอนเอียงไปทางบวกหรือลบ หรือเขียนได้ว่า $E(\varepsilon | X) = 0$ แนวคิดนี้สำคัญมาก เพราะหมายถึงแบบจำลองไม่ได้พยากรณ์สูงเกินจริงหรือต่ำเกินจริง หากสมมติฐานนี้ไม่จริง แสดงว่าอาจมีปัจจัยสำคัญบางอย่างตกหล่นจากแบบจำลอง หรือมีอคติในการกำหนดรูปแบบสมการ ส่งผลให้ค่าประมาณ β_0 และ β_1 มีความเอนเอียง ตัวอย่างเช่น หากกำลังพยากรณ์ยอดขายจากราคาเพียงตัวเดียว แต่จริง ๆ แล้วยังมีผลจากฤดูกาลหรือกิจกรรมส่งเสริมการขายที่ไม่ได้ใส่ในแบบจำลอง ค่าคลาดเคลื่อนอาจมีรูปแบบไม่เฉลี่ยเป็นศูนย์อย่างแท้จริง

สมมติฐานที่ 4 ความแปรปรวนคงที่ของค่าคลาดเคลื่อน (Homoscedasticity) หมายความว่า ค่าคลาดเคลื่อนควรมีความแปรปรวนเท่า ๆ กันในทุกระดับของ X หรือเขียนว่า $var(\varepsilon | X) = \sigma^2$ คงที่ หากความแปรปรวนของค่าคลาดเคลื่อนเพิ่มขึ้นหรือลดลงตามระดับของ X จะเรียกว่าเกิดปัญหา Heteroscedasticity ตัวอย่างเช่น ในข้อมูลรายได้กับการใช้จ่าย คนที่มีรายได้สูงอาจมีความกระจายของการใช้จ่ายมากกว่าคนรายได้ต่ำ หากปัญหานี้เกิดขึ้น ถึงแม้ค่าประมาณสัมประสิทธิ์ยังอาจไม่ลำเอียง แต่ค่าความคลาดเคลื่อนมาตรฐานที่คำนวณได้จะผิด ทำให้ค่า t-test หรือ F-test ผิดพลาด ส่งผลให้การทดสอบสมมติฐานคลาดเคลื่อนตาม นอกจากนี้ช่วงความเชื่อมั่นจะแคบหรือกว้างผิดจากความจริงทำให้ไม่น่าเชื่อถือ การตรวจสอบนิยมใช้กราฟค่าความคลาดเคลื่อน (residual plot) ถ้าจุดกระจายมีลักษณะบานออกเหมือนพัดหรือกรวย แสดงว่าความแปรปรวนไม่คงที่ ตัวอย่างปัญหาค่าคลาดเคลื่อนมีความแปรปรวนไม่คงที่ (Heteroscedasticity) แสดงดังภาพที่ 6.2

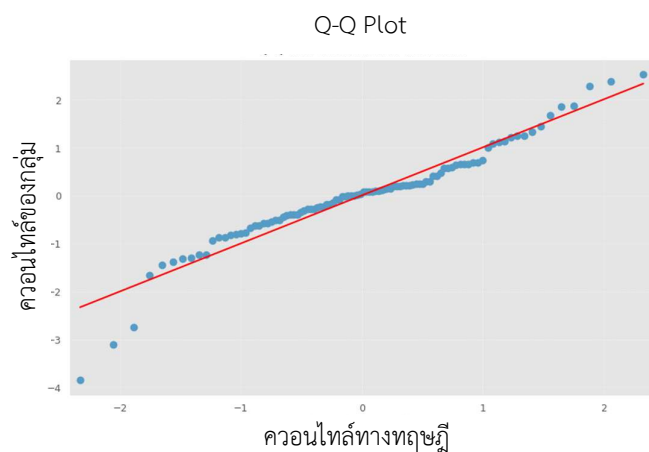


ภาพที่ 6.2 ตัวอย่างปัญหา Heteroscedasticity

ที่มา: ภาพโดยผู้แต่ง

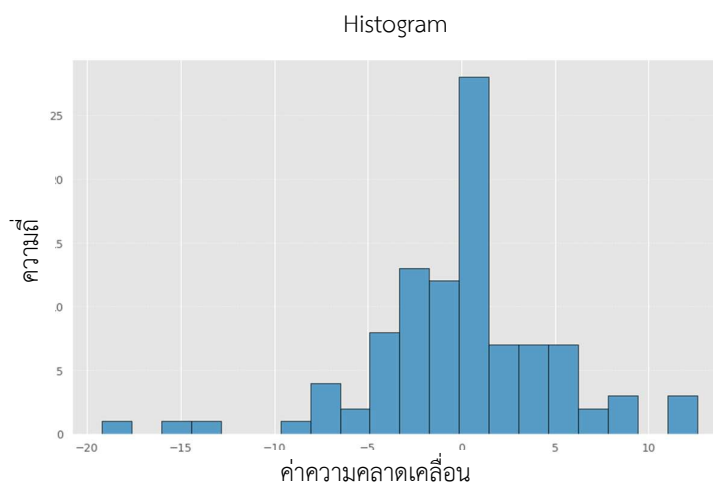
สมมติฐานที่ 5 คือ ค่าคลาดเคลื่อนมีการแจกแจงแบบปกติ (normality of errors) หมายถึง ค่าคลาดเคลื่อน ε ควรมีการแจกแจงแบบปกติด้วยค่าเฉลี่ยศูนย์และความแปรปรวนคงที่ หรือเขียนว่า $\varepsilon \sim N(0, \sigma^2)$ สมมติฐานนี้มีความสำคัญมากเมื่อเราต้องการใช้การอนุมานทางสถิติ เช่น การทดสอบนัยสำคัญของสัมประสิทธิ์หรือการสร้างช่วงความเชื่อมั่น ในกรณีตัวอย่างขนาดใหญ่ ผลกระทบของการไม่เป็นปกติอาจลดลงตามทฤษฎีขีดจำกัดส่วนกลาง (central theorem) แต่ถ้าตัวอย่างมีขนาดเล็ก การละเมิดสมมติฐานนี้อาจทำให้ผลทดสอบไม่น่าเชื่อถือ การตรวจสอบอาจใช้ Histogram ของ residual, Q-Q Plot หรือการทดสอบทางสถิติ เช่น Shapiro-Wilk test ตัวอย่างปัญหาค่าคลาดเคลื่อนมีการแจกแจงไม่เป็นปกติสามารถตรวจสอบโดยใช้ Q-Q Plot หรือ Histogram ในภาพที่ 6.3 และ 6.4

บทที่ 6 การพยากรณ์ความต้องการสินค้าด้วยวิธีการเรียนรู้ของเครื่อง



ภาพที่ 6.3 Q-Q Plot

ที่มา: ภาพโดยผู้แต่ง



ภาพที่ 6.4 Histogram

ที่มา: ภาพโดยผู้แต่ง

จากภาพที่ 6.3 สังเกตจาก Q-Q Plot พบว่า ค่าคลาดเคลื่อนไม่เป็นปกติเนื่องจากการกระจายไม่เกาะกับเส้นตรงอย่างมีนัยสำคัญ และเมื่อสังเกตจากภาพที่ 6.4 พบว่าค่าคลาดเคลื่อนมีการกระจายไม่สมมาตร (เบ้ซ้าย) แสดงถึงการแจกแจงที่ไม่เป็นปกติ (nonnormality)

6.1.2 การถดถอยเชิงเส้นพหุคูณ

การถดถอยเชิงเส้นพหุคูณ (Multiple Linear Regression, MLR) เป็นวิธีการทางสถิติที่ใช้ศึกษาความสัมพันธ์ระหว่างตัวแปรตาม (dependent Variable) 1 ตัวกับ ตัวแปรอิสระ (independent Variables) มากกว่า 1 ตัว เป็นการขยายจาก Simple Linear Regression ซึ่งมีตัวแปรอิสระเพียงตัวเดียว ให้สามารถวิเคราะห์ปัจจัยหลายตัวที่ส่งผลต่อตัวแปรตามตัวเดียวกันได้

สมการของ Multiple Regression

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon \quad (6.2)$$

โดยที่

Y = ตัวแปรตาม

$X_1, X_2, \dots, X_k$ = ตัวแปรอิสระหลายตัว

β_0 = ค่าคงที่ (intercept)

$\beta_1, \beta_2, \dots, \beta_k$ = ค่าสัมประสิทธิ์ของตัวแปรอิสระ

ε = ค่าคลาดเคลื่อน

ตัวอย่าง 6.1 การถดถอยเชิงเส้นพหุคูณ

การพยากรณ์อุปสงค์สินค้าขึ้นอยู่กับการปัจจัย ได้แก่ ราคา โปรโมชั่น รายได้ครัวเรือน เป็นต้น โดยทำการเก็บข้อมูลของปี ค.ศ. 2020 แสดงดังตารางที่ 6.1

ตารางที่ 6.1 ข้อมูลยอดขายและปัจจัยที่เกี่ยวข้องของปี ค.ศ. 2020

#	วันที่	ราคา (บาท)	โปรโมชั่น	รายได้ครัวเรือน (บาท)	ยอดขาย(Y) (หน่วย)
0	2020-01-01	10.618	0	32610.958	1064
1	2020-02-01	19.261	0	30142.510	715
2	2020-03-01	15.980	0	31564.911	882
3	2020-04-01	13.980	1	29680.974	997
4	2020-05-01	7.340	0	31050.343	1264
5	2020-06-01	7.340	0	30242.226	1313
6	2020-07-01	5.871	0	30797.007	1155
7	2020-08-01	17.993	0	31893.832	792
8	2020-09-01	14.017	0	29167.579	816

#	วันที่	ราคา (บาท)	โปรโมชั่น	รายได้ครัวเรือน (บาท)	ยอดขาย(Y) (หน่วย)
9	2020-10-01	15.621	0	35089.297	978
10	2020-11-01	5.309	0	28992.990	1238
11	2020-12-01	19.549	0	28677.150	416

ที่มา: จัดทำโดยผู้แต่ง

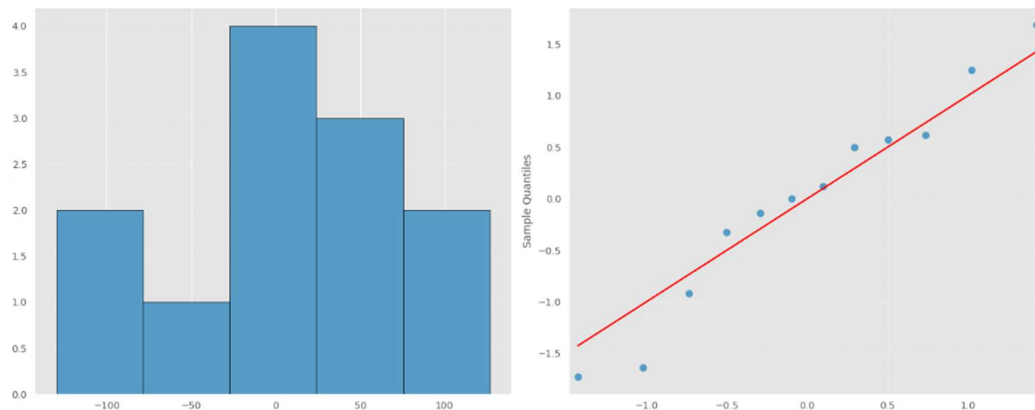
ในตัวอย่างนี้ ตัวแปรตามมี 1 ตัว คือ ยอดขาย ในขณะที่ตัวแปรอิสระมี 3 ตัวได้แก่ ราคา โปรโมชั่น และ รายได้ครัวเรือน โดยทำการเก็บข้อมูล 1 ปี การวิเคราะห์การถดถอยเชิงเส้นพหุคูณ สามารถทำได้โดยใช้โปรแกรมสำเร็จรูปอย่างเช่น Python ในการ fit โมเดล ในตัวอย่างนี้ใช้โปรแกรม ดังกล่าวในการ fit โมเดล โดยผลลัพธ์ที่ได้แสดงดังตารางที่ 6.2

ตารางที่ 6.2 ค่าสัมประสิทธิ์

Terms	β_i
Intercept	247.757
Price	-47.048
Promotion	146.801
Income	0.042

ที่มา: จัดทำโดยผู้แต่ง

หลังจากนั้นทำการทดสอบสมมติฐานว่าไม่ถูกละเมิดโดยการทำ Q-Q Plot และ Histogram แสดงดัง ภาพที่ 6.5



ภาพที่ 6.5 Q-Q Plot และ Histogram ของค่าคลาดเคลื่อน

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลจาก โปรแกรม Python)

จากภาพที่ 6.5 (กราฟด้านซ้าย) พบว่าฮิสโตแกรมนี้แสดงการกระจายความถี่ของค่าคลาดเคลื่อน หากค่าคลาดเคลื่อนมีการแจกแจงแบบปกติ กราฟควรมีลักษณะเป็นรูปโค้งระฆัง (bell-shaped curve) และมีความสมมาตรรอบค่าศูนย์ ฮิสโตแกรมของค่าคลาดเคลื่อนจากแบบจำลองปี 2020 แสดงการกระจายที่ค่อนข้างมีลักษณะคล้ายโค้งระฆังและมีจุดศูนย์กลางอยู่ใกล้ศูนย์ ซึ่งบ่งชี้ว่าค่าคลาดเคลื่อนมีการแจกแจงที่ใกล้เคียงกับการแจกแจงแบบปกติในระดับที่เหมาะสม

Q-Q Plot (กราฟด้านขวา) เป็นกราฟที่ใช้เปรียบเทียบควอนไทล์ของค่าคลาดเคลื่อนกับควอนไทล์ของการแจกแจงปกติทางทฤษฎี หากค่าคลาดเคลื่อนมีการแจกแจงแบบปกติ จุดข้อมูลควรเรียงตัวอยู่ใกล้กับเส้นตรงที่มีมุม 45 องศา สำหรับแบบจำลองปี 2020 จุดข้อมูลส่วนใหญ่ใน Q-Q Plot เรียงตัวใกล้กับเส้นตรง โดยเฉพาะบริเวณกึ่งกลางของการแจกแจง แม้ว่าอาจมีการเบี่ยงเบนเล็กน้อยที่บริเวณปลายของการแจกแจง (tail) แต่โดยรวมแล้วแสดงให้เห็นว่าค่าคลาดเคลื่อนมีการแจกแจงที่ใกล้เคียงกับการแจกแจงแบบปกติ จากตารางที่ 6.2 เราสามารถเขียนสมการถดถอยเชิงเส้นพหุคูณแสดงดังสมการที่ 6.3

$$Y = 247.757 - 47.048(\text{Price}) + 146.801(\text{Promotion}) + 0.042(\text{Income}) \quad (6.3)$$

ตัวอย่างการใช้สมการที่ได้ยกตัวอย่างเช่น กรณีที่ไม่มี Promotion (Promotion=0) สมการจะลดรูปไปเป็นสมการที่ 6.4

$$Y = 247.757 - 47.048(\text{Price}) + 0.042(\text{Income}) \quad (6.4)$$

ถ้าสมมติ ราคา (Price) = 15 บาท และ รายได้ครัวเรือน = 30,000 บาท เราสามารถพยากรณ์ ยอดขายได้ดังต่อไปนี้

$$\begin{aligned} Y &= 247.757 - 47.048(\text{Price}) + 0.042(\text{Income}) \\ Y &= 247.757 - 47.048(15) + 0.042(30,000) \approx 2,213 \end{aligned}$$

แต่ถ้ามีการทำโปรโมชันดังสมการที่ 6.3 ยอดขายจะเพิ่มขึ้น จากเดิมถ้าคงราคาและรายได้ครัวเรือนไว้ที่ค่าเดิม แสดงดังต่อไปนี้

$$Y = 247.757 - 47.048(15) + 146.801(1) + 0.042(30,000) \approx 2,360$$

6.2 แบบจำลองป่าสุ่ม (Random Forest, RF)

ในอดีต การพยากรณ์ความต้องการหรืออุปสงค์สินค้ามักอาศัยแบบจำลองทางสถิติ เช่น Exponential Smoothing และ ARIMA อย่างไรก็ตาม เมื่อข้อมูลทางธุรกิจมีปริมาณมากขึ้นและมีปัจจัยที่ส่งผลต่ออุปสงค์จำนวนมาก วิธีการทางสถิติแบบดั้งเดิมอาจไม่สามารถอธิบายความสัมพันธ์ที่ซับซ้อนระหว่างตัวแปรได้อย่างมีประสิทธิภาพ ด้วยเหตุนี้ เทคนิค ML จึงถูกนำมาใช้มากขึ้นในการพยากรณ์อุปสงค์ในซัพพลายเชน เนื่องจากสามารถเรียนรู้รูปแบบข้อมูลที่ไม่เป็นเชิงเส้นและมีปฏิสัมพันธ์ระหว่างตัวแปรหลายตัวได้ดี (Carbonneau et al., 2008; Makridakis et al., 2018) หนึ่งในอัลกอริทึมที่ได้รับความนิยมในการพยากรณ์ความต้องการสินค้าคือ Random Forest ซึ่งเป็นวิธีการเรียนรู้แบบ ensemble learning ที่พัฒนาโดย Breiman (2001) โดยอาศัยแนวคิดของการรวมผลลัพธ์จากต้นไม้ตัดสินใจจำนวนมากเพื่อเพิ่มความแม่นยำของการพยากรณ์ วิธีการนี้ใช้การสุ่มตัวอย่างข้อมูลแบบ Bootstrap และการสุ่มเลือกตัวแปรในการสร้างต้นไม้แต่ละต้น ซึ่งช่วยลดความแปรปรวนของแบบจำลองและลดปัญหา overfitting เมื่อเทียบกับการใช้ต้นไม้ตัดสินใจเพียงต้นเดียว (Breiman, 2001; Schonlau & Zou, 2020)

หลักการทำงานของ RF คือการสร้างต้นไม้ตัดสินใจจำนวนมากจากชุดข้อมูลที่สุ่มขึ้นมา และนำผลลัพธ์จากต้นไม้ทั้งหมดมารวมกันเพื่อให้ได้ค่าพยากรณ์สุดท้าย ในกรณีของการพยากรณ์เชิงตัวเลข เช่น การพยากรณ์ความต้องการสินค้า ค่าพยากรณ์สุดท้ายจะได้จากค่าเฉลี่ยของผลลัพธ์จาก

ต้นไม้ทุกต้น แนวทางนี้ช่วยให้แบบจำลองสามารถเรียนรู้รูปแบบข้อมูลที่ซับซ้อนและมีความไม่เป็นเชิงเส้นได้ดี รวมทั้งมีความทนทานต่อข้อมูลที่มีความผันผวนสูง (Scornet et al., 2015; Genuer et al., 2015) ในบริบทของซัพพลายเชน RF สามารถนำมาใช้พยากรณ์ความต้องการสินค้าโดยใช้ข้อมูลจากหลายแหล่ง เช่น ยอดขายย้อนหลัง ราคาสินค้า โปรโมชัน วันหยุดทางการค้า ฤดูกาล สภาพอากาศ และข้อมูลพฤติกรรมลูกค้า การใช้ข้อมูลหลายมิติช่วยให้แบบจำลองสามารถเรียนรู้ความสัมพันธ์ระหว่างปัจจัยต่าง ๆ ที่ส่งผลต่ออุปสงค์สินค้าได้อย่างมีประสิทธิภาพ งานวิจัยหลายชิ้นพบว่า RF สามารถให้ผลการพยากรณ์ที่มีความแม่นยำสูงในงานค้าปลีกและการจัดการสินค้าคงคลัง โดยเฉพาะในกรณีที่ข้อมูลมีความสัมพันธ์ที่ซับซ้อนและไม่เป็นเชิงเส้น (Islam & Amin, 2020; Mitra et al., 2022)

นอกจากนี้ RF ยังถูกนำมาใช้ในการพยากรณ์ความต้องการสินค้าใหม่ในซัพพลายเชน โดยใช้ข้อมูลยอดขายของสินค้าในกลุ่มเดียวกันเพื่อเรียนรู้รูปแบบอุปสงค์ ตัวอย่างเช่น งานวิจัยของ Van Steenberghe et al. (2020) ได้นำเสนอวิธีการที่เรียกว่า DemandForest ซึ่งใช้ RF ในการเรียนรู้รูปแบบอุปสงค์ของสินค้าใหม่ โดยอาศัยข้อมูลจากสินค้าเก่าที่มีลักษณะคล้ายคลึงกัน วิธีการดังกล่าวสามารถช่วยให้องค์กรสามารถวางแผนการผลิตและการจัดการสินค้าคงคลังสำหรับสินค้าใหม่ได้อย่างมีประสิทธิภาพ การใช้ RF ในการพยากรณ์ความต้องการสินค้ายังได้รับการประยุกต์ใช้ในธุรกิจค้าปลีกขนาดใหญ่ ซึ่งมีข้อมูลยอดขายจำนวนมากในหลายช่องทางการจำหน่าย งานวิจัยหลายชิ้นพบว่าแบบจำลอง ML เช่น RF และ ANN สามารถปรับปรุงความแม่นยำของการพยากรณ์ยอดขายได้เมื่อเทียบกับวิธีการแบบดั้งเดิม โดยเฉพาะในกรณีที่ข้อมูลขนาดใหญ่และมีรูปแบบความต้องการที่ซับซ้อน (Bandara et al., 2020; Mediavilla et al., 2022) ข้อดีสำคัญของ RF ในการพยากรณ์อุปสงค์คือความสามารถในการจัดการกับข้อมูลจำนวนมาก ความสัมพันธ์ที่ไม่เป็นเชิงเส้น และความไม่แน่นอนของข้อมูล นอกจากนี้ยังสามารถวิเคราะห์ความสำคัญของตัวแปร (feature importance) ได้ ทำให้ผู้วิเคราะห์สามารถระบุปัจจัยที่มีผลต่ออุปสงค์สินค้าได้อย่างชัดเจน ซึ่งมีประโยชน์ต่อการวางแผนกลยุทธ์ทางธุรกิจและการบริหารซัพพลายเชน อย่างไรก็ตาม RF ก็มีข้อจำกัดบางประการ เช่น ความยากในการตีความโมเดลเมื่อเทียบกับ Regression แบบดั้งเดิม และความต้องการทรัพยากรคอมพิวเตอร์ที่สูงขึ้นเมื่อใช้กับข้อมูลขนาดใหญ่

โดยสรุป RF เป็นเครื่องมือที่มีประสิทธิภาพสำหรับการพยากรณ์อุปสงค์ในซัพพลายเชน เนื่องจากสามารถเรียนรู้รูปแบบข้อมูลที่ซับซ้อนและปรับตัวเข้ากับข้อมูลขนาดใหญ่ได้ดี การนำเทคนิค

นี้มาใช้ร่วมกับข้อมูลธุรกิจที่หลากหลายสามารถช่วยเพิ่มความแม่นยำของการพยากรณ์ และสนับสนุนการตัดสินใจด้านการวางแผนการผลิต การจัดซื้อ และการจัดการสินค้าคงคลังในองค์กรได้อย่างมีประสิทธิภาพ

6.2.1 ขั้นตอนการทำ Random Forest (RF)

การพัฒนาแบบจำลอง RF เพื่อพยากรณ์อุปสงค์สินค้าในซัพพลายเชนโดยทั่วไปจะมีขั้นตอนดังต่อไปนี้

1. การกำหนดปัญหาและตัวแปรเป้าหมาย โดยสำหรับการพยากรณ์อุปสงค์ ตัวแปรเป้าหมายจะเป็น ความต้องการสินค้า หรือ ยอดขายของสินค้านั้น
2. การรวบรวมข้อมูล เช่น ข้อมูลยอดขายย้อนหลัง ราคาสินค้า โปรโมชั่น วันหยุด สภาพอากาศ และอื่นๆที่เป็น
3. การเตรียมและทำความสะอาดข้อมูล ก่อนจะไปยังขั้นตอนการสร้างโมเดลพยากรณ์จำเป็นที่จะต้องจัดการข้อมูลให้ถูกต้องเพื่อความแม่นยำของโมเดล ข้อมูลที่ต้องจัดการอย่างเช่น ค่าที่หายไป (missing values) ข้อมูลที่ผิดปกติ (outliers) การแปลงรูปแบบข้อมูลทางเวลาเช่น แปลงจากเดือนเป็นฤดูกาลเพื่อให้สอดคล้องกับช่วงเวลาที่ต้องการพยากรณ์
4. การแบ่งชุดข้อมูล โดยข้อมูลจะถูกแบ่งออกเป็น 2 ส่วนได้แก่ ชุดฝึก (training) กับ ชุดทดสอบ (test) โดยชุดฝึกจะมีวัตถุประสงค์ไว้เพื่อสร้างโมเดล ในขณะที่ชุดทดสอบไว้ใช้ประเมินความเที่ยงตรงของโมเดล
5. การสร้างโมเดล RF มีขั้นตอนโดยสรุปเริ่มจากการสุ่มตัวอย่างข้อมูลจากชุดฝึกด้วยวิธี Bootstrapping หลังจากนั้นทำการสร้างต้นไม้ตัดสินใจจากข้อมูลที่สุ่มมา โดยในแต่ละจุดแตกกิ่งของต้นไม้ จะสุ่มเลือกบางตัวแปรมาใช้ในการแบ่งข้อมูล แล้วจึงทำซ้ำหลายๆครั้งจนได้ต้นไม้จำนวนมาก
6. การปรับพารามิเตอร์ของโมเดล จะต้องทำการปรับค่าของพารามิเตอร์ที่สำคัญหลายตัว ได้แก่ จำนวนต้นไม้ ความลึกสูงสุดของต้นไม้ จำนวนตัวแปรที่ใช้ในแต่ละการแบ่งกิ่ง จำนวนตัวอย่างขั้นต่ำที่ใช้แบ่งโหนด โดยสามารถใช้วิธีอย่างเช่น grid search หรือ random search ในการปรับค่าที่ให้ผลพยากรณ์ที่ดีที่สุด
7. การพยากรณ์ผลลัพธ์ โดยผลลัพธ์ของการพยากรณ์จะเป็นค่าเฉลี่ยที่เกิดจากผลลัพธ์ของต้นไม้ทุกๆต้น
8. การประเมินความเที่ยงตรงสามารถใช้ตัวชี้วัดมาตรฐานหลายๆตัวประกอบกันในการประเมินความเที่ยงตรง ได้แก่ MAE RMSE MAPE เป็นต้น

9. การวิเคราะห์ความสำคัญของปัจจัยหรือตัวแปร เนื่องจากวิธี RF สามารถจะจัดลำดับความสำคัญของปัจจัยที่เกี่ยวข้องแต่ละตัวเพื่อดูว่าปัจจัยใดที่มีผลกระทบต่ออุปสงค์มากที่สุด เพื่อนำไปใช้ต่อยอดในการวางแผนทางการตลาด (Molnar, 2022)

10. การนำไปใช้งานจริงโดยหลังจากที่แบบจำลองมีความเที่ยงตรงในระดับที่ยอมรับได้อย่างมีนัยสำคัญทางสถิติแล้ว เราสามารถนำโมเดลที่พัฒนาขึ้นไปใช้งานจริงได้ (Molnar, 2022)

ตัวอย่าง 6.2 แบบจำลองป่าสุ่ม

หมายเหตุ: ตัวอย่างนี้ใช้ข้อมูลอย่างง่ายเพื่อให้สามารถทำการคำนวณด้วยมือและทำความเข้าใจเบื้องต้นได้ แต่สำหรับข้อมูลจริงๆจะมีขนาดใหญ่และยากที่จะคำนวณด้วยมือจำเป็นที่จะต้องใช้โปรแกรมอย่างเช่น PYTHON ในการช่วยวิเคราะห์ ซึ่งจะแสดงรายละเอียดในบทการนำไปใช้งานจริง

สมมติว่าบริษัทต้องการพยากรณ์ ยอดขายเครื่องดื่มโดยใช้ตัวแปร 2 ตัวแปรได้แก่ โปรโมชัน ($1 =$ มีโปรโมชัน, $0 =$ ไม่มี) และ อุณหภูมิ โดยได้ทำการเก็บข้อมูลเบื้องต้นและได้ข้อมูลชุดฝึก มาเพื่อทำการสร้างโมเดลพยากรณ์ ข้อมูลชุดฝึกแสดงดังตารางที่ 6.3 โดยผู้จัดการบริษัทต้องการพยากรณ์ว่า ถ้าจัดให้มีโปรโมชันแล้ว อุณหภูมิ 30 องศา จะมียอดขายเท่าไร?

ตารางที่ 6.3 ข้อมูลชุดฝึก

วัน	โปรโมชัน	อุณหภูมิ	ยอดขาย
1	1	30	120
2	0	28	90
3	1	31	130
4	0	27	85
5	0	29	100
6	1	32	140

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 1 การสุ่มตัวอย่างด้วยวิธี Bootstrapping

วิธี RF จะทำการสุ่มตัวอย่างข้อมูลเพื่อทำการสร้างต้นไม้ตัดสินใจหลายๆต้นในตัวอย่างนี้จะสมมติให้เห็นการทำงานเบื้องต้นจึงทำการสุ่มตัวอย่างมาเพียง 3 ต้น แสดงดังตารางที่ 6.4 -6.6

ตารางที่ 6.4 ข้อมูลสุ่มตัวอย่างของต้นไม้ต้นที่ 1

โปรโมชัน	1	1	1	0	0	1
อุณหภูมิ	30	31	31	27	29	32
ยอดขาย	120	130	130	85	100	140

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 6.5 ข้อมูลสุ่มตัวอย่างของต้นไม้ต้นที่ 2

โปรโมชัน	0	0	1	0	1	1
อุณหภูมิ	28	28	31	29	32	32
ยอดขาย	90	90	130	100	140	140

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 6.6 ข้อมูลสุ่มตัวอย่างของต้นไม้ต้นที่ 3

โปรโมชัน	1	0	0	0	0	1
อุณหภูมิ	30	28	27	27	29	32
ยอดขาย	120	90	85	85	100	140

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 2 สร้างต้นไม้ตัดสินใจ เพื่อความง่ายในการคำนวณ จะสมมติให้ต้นไม้แตกกิ่งจากตัวแปรเดียว

ต้นไม้ตัดสินใจ 1 แยกที่โปรโมชัน

โปรโมชัน = 1

ถ้าโปรโมชัน = 1; ยอดขาย = 120, 130, 130, 140;

ค่าเฉลี่ย = $\frac{(120+130+130+140)}{4} = 130$

โปรโมชัน = 0

ถ้าโปรโมชัน = 0; ยอดขาย = 85, 100;

$$\text{ค่าเฉลี่ย} = \frac{(85+1)}{2} = 92.5$$

ต้นไม้ตัดสินใจ 2 แยกที่อุณหภูมิ

$$\text{อุณหภูมิ} \geq 30$$

ถ้าอุณหภูมิ ≥ 30 ; ยอดขาย = 130, 140, 140;

$$\text{ค่าเฉลี่ย} = \frac{(130+140+140)}{3} = 136.67$$

$$\text{อุณหภูมิ} < 30$$

ถ้าอุณหภูมิ < 30 ; ยอดขาย = 90, 90, 100;

$$\text{ค่าเฉลี่ย} = \frac{(90+90+100)}{3} = 93.33$$

ต้นไม้ตัดสินใจ 3 แยกที่โปรโมชั่น

$$\text{โปรโมชั่น} = 1$$

ถ้าโปรโมชั่น = 1; ยอดขาย = 120, 140;

$$\text{ค่าเฉลี่ย} = \frac{(120+140)}{2} = 130$$

$$\text{โปรโมชั่น} = 0$$

ถ้าโปรโมชั่น = 0; ยอดขาย = 90, 85, 85, 100;

$$\text{ค่าเฉลี่ย} = \frac{(90+85+85+10)}{4} = 90$$

สรุป ถ้าผู้จัดการต้องการพยากรณ์ว่าถ้าจัดให้มีโปรโมชั่นแล้ว อุณหภูมิ 30 องศา จะมี ยอดขายเท่าไร เราจะได้ว่า (โปรโมชั่น=1 อุณหภูมิ=30) จากต้นไม้ตัดสินใจต้นที่ 1 จะได้ว่า ถ้า โปรโมชั่น = 1 ยอดขาย = 130 จากต้นไม้ตัดสินใจต้นที่ 2 จะได้ว่า อุณหภูมิ = 30 ยอดขาย = 136.67 และสุดท้าย จากต้นไม้ตัดสินใจต้นที่ 3 ถ้า โปรโมชั่น = 1 ยอดขาย = 130 ดังนั้นเราต้อง ดำเนินการต่อในขั้นตอนที่ 3

ขั้นตอนที่ 3 รวมผล RF

RF จะใช้ค่าเฉลี่ยแสดงดังต่อไปนี้

$$\hat{y} = \frac{130+1 \cdot 67+130}{3} = 132.22$$

ดังนั้นเราสามารถพยากรณ์ยอดขายกรณี (โปรโมชัน=1 อุณหภูมิ=30) จะได้เท่ากับ 132.22 หน่วย

6.3 แบบจำลองโครงข่ายประสาทเทียม (Artificial Neural Network, ANN)

แบบจำลอง โครงข่ายประสาทเทียม (ANN) เป็นหนึ่งในเทคนิคการเรียนรู้ของเครื่องที่ได้รับ ความนิยมอย่างมากในการพยากรณ์ อุปสงค์ เนื่องจากมีความสามารถในการเรียนรู้รูปแบบ ความสัมพันธ์ที่ซับซ้อนและไม่เป็นเชิงเส้นระหว่างตัวแปรต่าง ๆ ได้ดีกว่าวิธีการทางสถิติแบบดั้งเดิม เช่น การถดถอยเชิงเส้นหรือแบบจำลองอนุกรมเวลา ในงานด้านการพยากรณ์ธุรกิจและเศรษฐศาสตร์ พบว่า ANN สามารถจำลองรูปแบบข้อมูลที่มีความซับซ้อนได้อย่างมีประสิทธิภาพ และให้ความ แม่นยำในการพยากรณ์ที่ดีในหลายบริบท (Zhang, Patuwo, & Hu, 1998) ในบริบทของการจัดการ ชัพพลายเชน เทคโนโลยีการวิเคราะห์ข้อมูลขนาดใหญ่และการเรียนรู้ของเครื่องทำให้แบบจำลองที่ใช้ โครงข่ายประสาทเทียมถูกนำมาใช้ในการพยากรณ์อุปสงค์ การจัดการสินค้าคงคลัง และการวางแผนการผลิตมากขึ้น (Choi, Wallace, & Wang, 2018)

โครงข่ายประสาทเทียมได้รับแรงบันดาลใจจากโครงสร้างของระบบประสาทในสมองมนุษย์ โดยประกอบด้วยหน่วยประมวลผลที่เรียกว่า นิวรอน (Neurons) ซึ่งเชื่อมต่อกันเป็นเครือข่ายหลาย ชั้น ได้แก่ ชั้นข้อมูลนำเข้า (input layer) ชั้นซ่อน (hidden layer) และ ชั้นผลลัพธ์ (output layer) ข้อมูลจะถูกส่งผ่านเครือข่ายโดยมีการคำนวณผลรวมถ่วงน้ำหนักของตัวแปรนำเข้า และส่งผ่าน ฟังก์ชันกระตุ้น (activation function) เพื่อสร้างค่าผลลัพธ์ของแบบจำลอง กระบวนการเรียนรู้ของ เครือข่ายจะปรับค่าถ่วงน้ำหนักผ่านอัลกอริทึม เช่น backpropagation เพื่อให้ค่าคลาดเคลื่อน ระหว่างค่าพยากรณ์กับค่าจริงต่ำที่สุด

ในการพยากรณ์อุปสงค์ของซัพพลายเชน ตัวแปรนำเข้าอาจประกอบด้วยข้อมูลยอดขายใน อดีต ระดับราคา กิจกรรมส่งเสริมการขาย ฤดูกาล หรือปัจจัยทางเศรษฐกิจ ในขณะที่ตัวแปรผลลัพธ์ คือปริมาณอุปสงค์ในอนาคต การใช้ ANN ทำให้สามารถจับรูปแบบความสัมพันธ์ที่ซับซ้อนของปัจจัย เหล่านี้ได้ ซึ่งช่วยเพิ่มความแม่นยำของการพยากรณ์เมื่อเทียบกับวิธีการพยากรณ์แบบดั้งเดิม งานวิจัย ในด้านการจัดการซัพพลายเชนพบว่าเทคนิคการเรียนรู้ของเครื่อง เช่น Neural Networks สามารถ ปรับปรุงประสิทธิภาพของการพยากรณ์อุปสงค์ได้ โดยเฉพาะในกรณีที่ข้อมูลมีความไม่เป็นเชิงเส้น หรือมีความผันผวนสูง (Carbonneau, Laframboise, & Vahidov, 2008)

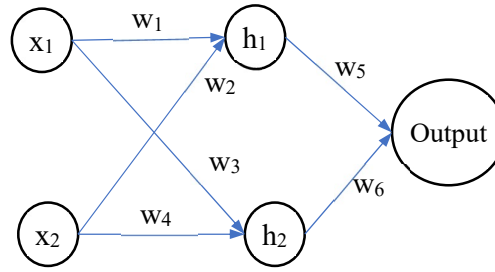
นอกจากนี้ งานวิจัยด้านการพยากรณ์ในภาคค้าปลีกยังชี้ให้เห็นว่าการใช้วิธีการขั้นสูง เช่น ML และ ANN มีบทบาทสำคัญในการพยากรณ์ความต้องการของลูกค้าในสภาพแวดล้อมธุรกิจที่มีความซับซ้อนมากขึ้น ขณะเดียวกัน การศึกษาด้านการพยากรณ์ระดับสากล เช่น การแข่งขัน M4 forecasting competition ยังแสดงให้เห็นว่าการพัฒนาเทคนิคการพยากรณ์สมัยใหม่ รวมถึงวิธีการที่ใช้ ML มีศักยภาพในการเพิ่มความแม่นยำของแบบจำลองพยากรณ์อย่างมีนัยสำคัญ (Makridakis, Spiliotis, & Assimakopoulos, 2018) นอกจากนี้ การพัฒนาแบบจำลองพยากรณ์สมัยใหม่ยังรวมถึงการผสมผสานข้อมูลหลายความถี่และโครงสร้างของอนุกรมเวลา ซึ่งช่วยเพิ่มประสิทธิภาพของแบบจำลองพยากรณ์ในสถานการณ์ที่มีข้อมูลซับซ้อน

แม้ว่าโครงข่ายประสาทเทียมจะมีข้อดีในด้านความสามารถในการเรียนรู้รูปแบบข้อมูลที่ซับซ้อนและความแม่นยำของการพยากรณ์ แต่การใช้งานในทางปฏิบัติก็ยังมีข้อจำกัดบางประการ เช่น ความต้องการข้อมูลจำนวนมากในการฝึกแบบจำลอง ความซับซ้อนในการกำหนดโครงสร้างของโครงข่าย และความยากในการตีความผลลัพธ์ของแบบจำลองเมื่อเทียบกับวิธีการทางสถิติแบบดั้งเดิม อย่างไรก็ตาม ด้วยความก้าวหน้าของเทคโนโลยีข้อมูลและการคำนวณ โครงข่ายประสาทเทียมยังคงเป็นหนึ่งในเครื่องมือสำคัญสำหรับการพยากรณ์อุปสงค์และการตัดสินใจเชิงวิเคราะห์ในระบบซัพพลายเชนสมัยใหม่

ตัวอย่าง 6.3 การจำลองโครงข่ายประสาทเทียม

หมายเหตุ: ตัวอย่างนี้ใช้ข้อมูลอย่างง่ายเพื่อให้สามารถทำการคำนวณด้วยมือและทำความเข้าใจเบื้องต้นได้ แต่สำหรับข้อมูลจริงๆจะมีขนาดใหญ่และยากที่จะคำนวณด้วยมือจำเป็นที่จะต้องใช้อุปกรณ์อย่างเช่น PYTHON ในการช่วยวิเคราะห์ ซึ่งจะแสดงรายละเอียดในบทการนำไปใช้งานจริง

สมมติว่าต้องการพยากรณ์ยอดขายของสินค้าชนิดหนึ่ง โดยมีตัวแปรนำเข้า 2 ตัว ได้แก่ x_1 = ราคา x_2 = ค่าใช้จ่ายในการทำโปรโมชั่น และ Y = ยอดขาย และ กำหนดให้ ชั้นนำเข้า (input Layer: x_1, x_2) มี 2 ตัวแปร ชั้นซ่อน (hidden Layer: h_1, h_2) มี 2 โหนด และ ชั้นผลลัพธ์ (output Layer) มี 1 โหนด แบบจำลองแสดงดังภาพที่ 6.6 ขั้นตอนพร้อมทั้งคำอธิบายของแต่ละขั้นตอนสำหรับการสร้างโมเดลแสดงดังต่อไปนี้



ภาพที่ 6.6 แบบจำลอง ANN ของตัวอย่างที่ 6.3

ที่มา: ภาพโดยผู้แต่ง

ขั้นตอนที่ 1 กำหนดข้อมูลนำเข้าของเดือนใหม่ ราคา (x_1) = 4 โปรโมชัน (x_2) = 3

ขั้นตอนที่ 2 กำหนดค่าน้ำหนักและค่าอคติเริ่มต้น แสดงดังตารางที่ 6.7

ตารางที่ 6.7 ค่าน้ำหนักและค่าอคติเริ่มต้น

	h_1	h_2		Output
x_1	$w_1 = 0.2$	$w_3 = 0.3$	h_1	$w_5 = 0.5$
x_2	$w_2 = 0.4$	$w_4 = 0.1$	h_2	$w_6 = 0.6$
b(Bias)	$b_1 = 0.1$	$b_2 = 0.2$	b(Bias)	$b_3 = 0.1$

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 3 ในขั้นตอนนี้จะเริ่มกระบวนการคำนวณส่งผ่านไปข้างหน้า (forward propagation) โดยจะเริ่มทำการคำนวณจาก ชั้นซ่อนก่อนโดยใช้สมการที่ 6.5-6.6

$$Z_1 = [W_{hidden}] \times [x] + [B_{hidden}] \quad (6.5)$$

$$A(Z_1) = \frac{1}{(1+e^{-Z_1})} \quad (6.6)$$

$$[W_{hidden}] = \begin{bmatrix} w_1 & w_2 \\ w_3 & w_4 \end{bmatrix} = \begin{bmatrix} 0.2 & 0.4 \\ 0.3 & 0.1 \end{bmatrix}$$

$$[x] = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 4 \\ 3 \end{bmatrix}$$

$$[B_{hidden}] = \begin{bmatrix} 0.1 \\ 0.2 \end{bmatrix}$$

$$Z_1 = \begin{bmatrix} 0.2 & 0.4 \\ 0.3 & 0.1 \end{bmatrix} \times \begin{bmatrix} 4 \\ 3 \end{bmatrix} + \begin{bmatrix} 0.1 \\ 0.2 \end{bmatrix} = \begin{bmatrix} 2.1 \\ 1.7 \end{bmatrix}$$

จากนั้นผ่าน Sigmoid function ในสมการที่ 7.6 ดังต่อไปนี้

$$A(Z_1) = \begin{bmatrix} \frac{1}{(1+e^{-2.1})} \\ \frac{1}{(1+e^{-1.7})} \end{bmatrix} = \begin{bmatrix} 0.891 \\ 0.845 \end{bmatrix}$$

ขั้นตอนที่ 4 ชั้นผลลัพธ์

$$[W_{output}] = [w_5 \quad w_6] = [0.5 \quad 0.6]$$

$$[B_{output}] = [b_3] = [0.1]$$

$$Z_2 = [0.5 \quad 0.6] \times \begin{bmatrix} 0.891 \\ 0.845 \end{bmatrix} + [0.1] = [1.0525]$$

จากนั้นผ่าน Sigmoid function จะได้ว่า

$$\hat{y} = A(Z_2) = \frac{1}{(1 + e^{-1.0525})} = 0.741$$

ถ้าทำการแปลงกลับจะได้ว่า $0.741 \times 1000 = 741$ หรือยอดขายเท่ากับ 741 หน่วย

ขั้นตอนที่ 5 จะเป็นขั้นตอนการปรับน้ำหนักโดยใช้การทำย้อนกลับ (back propagation) เพื่อให้ค่า error ลดลงและเข้าสู่ steady state โดยในตัวอย่างนี้จะแสดงเพียงแค่ 1 รอบ (iteration) เพื่อเป็นแนวทางการทำความเข้าใจในกระบวนการของวิธี ANN โดยสมมติให้ค่าจริง (y) เท่ากับ 0.800 ในขั้นตอนนี้จะเริ่มด้วยการหา error term ของชั้นผลลัพธ์โดยใช้สมการ 6.7

$$\delta^{(2)} = (A(Z_2) - y) \odot \dot{A}(Z_2) \quad (6.7)$$

$$\dot{A}(Z_2) = A(Z_2) \odot (1 - A(Z_2)) = 0.741(1 - 0.741) = 0.1919$$

ดังนั้น

$$\delta^{(2)} = (0.741 - 0.800) \odot 0.1919 = -0.0113$$

ที่ Output Layer หลังจากนั้น ต้องทำการหาค่า Gradient ของน้ำหนักในชั้นผลลัพธ์ แสดงดังต่อไปนี้

$$\frac{\partial E}{\partial W_{output}} = \delta^{(2)}(A(Z_1))^T = -0.0113[0.891 \quad 0.845] = [-0.0101 \quad -0.0095]$$

และ Gradient ของอคติคำนวณดังต่อไปนี้

$$\frac{\partial E}{\partial B_{output}} = \delta^{(2)} = -0.0113$$

ที่ Hidden Layer ทำการส่งความคลาดเคลื่อนไปยังชั้นซ่อนใช้สมการที่ 7.8

$$\delta^{(1)} = \delta^{(2)}(W_{output})^T \odot \dot{A}(Z_1) \quad (6.8)$$

$$\delta^{(1)} = (-0.0113) \begin{bmatrix} 0.5 \\ 0.6 \end{bmatrix} \odot \dot{A}(Z_1)$$

ต้องทำการคำนวณค่าอนุพันธ์ของ Sigmoid ที่ชั้นซ่อน $\{\dot{A}(Z_1)\}$ ดังต่อไปนี้

$$\dot{A}(Z_1) = A(Z_1) \odot (1 - A(Z_1)) = \begin{bmatrix} 0.891 \\ 0.845 \end{bmatrix} \odot \begin{bmatrix} 1 - 0.891 \\ 1 - 0.845 \end{bmatrix} = \begin{bmatrix} 0.891(1 - 0.891) \\ 0.845(1 - 0.845) \end{bmatrix}$$

จะได้ว่า

$$\dot{A}(Z_1) = \begin{bmatrix} 0.0971 \\ 0.1310 \end{bmatrix}$$

คำนวณค่าความคลาดเคลื่อนที่ชั้นซ่อน

$$\delta^{(1)} = (-0.0113) \begin{bmatrix} 0.5 \\ 0.6 \end{bmatrix} \odot \begin{bmatrix} 0.0971 \\ 0.1310 \end{bmatrix} = \begin{bmatrix} -0.00565 \\ -0.00678 \end{bmatrix} \odot \begin{bmatrix} 0.0971 \\ 0.1310 \end{bmatrix} = \begin{bmatrix} -0.000549 \\ -0.000888 \end{bmatrix}$$

ขั้นตอนถัดมาคือการคำนวณหาค่า gradient ของน้ำหนักในชั้นซ่อน แสดงดังต่อไปนี้

$$\frac{\partial E}{\partial W_{hidden}} = \delta^{(1)}(x)^T = \begin{bmatrix} -0.000549 \\ -0.000888 \end{bmatrix} \begin{bmatrix} 4 & 3 \end{bmatrix} = \begin{bmatrix} -0.002196 & -0.001647 \\ -0.003552 & -0.002664 \end{bmatrix}$$

และการคำนวณหาค่า gradient ของอคติในชั้นซ่อน แสดงดังต่อไปนี้

$$\frac{\partial E}{\partial W_{hi}} = \delta^{(1)} = \begin{bmatrix} -0.000549 \\ -0.000888 \end{bmatrix}$$

หลังจากนั้นจะทำการปรับค่าน้ำหนักและอคติ โดยใช้สมการที่ 6.9-6.10

$$W_l^n = W_l - \eta \frac{\partial E}{\partial W_l} \quad (6.9)$$

$$B_l^n = B_l - \eta \frac{\partial E}{\partial B_l} \quad (6.10)$$

โดยที่ η = อัตราการเรียนรู้ ในตัวอย่างนี้กำหนดให้เท่ากับ 0.5

ที่ Output Layer ทำการอัปเดตชั้นผลลัพธ์

$$W_{output}^n = W_{output} - \eta \frac{\partial E}{\partial W_{output}} = \begin{bmatrix} 0.5 & 0.6 \end{bmatrix} - 0.5 \begin{bmatrix} -0.0101 & -0.0095 \end{bmatrix}$$

$$W_{output}^n \approx \begin{bmatrix} 0.5051 & 0.6048 \end{bmatrix}$$

ทำการอัปเดตค่าอคติเช่นเดียวกันจะได้

$$B_{output}^n = B_{output} - \eta \frac{\partial E}{\partial B_{output}} = 0.1 - 0.5(-0.0113) = 0.1057$$

ที่ Hidden Layer ทำการอัปเดตชั้นซ่อน

$$W_{hidden}^n = W_{hidden} - \eta \frac{\partial E}{\partial W_{hidden}} = \begin{bmatrix} 0.2 & 0.4 \\ 0.3 & 0.1 \end{bmatrix} - 0.5 \begin{bmatrix} -0.002196 & -0.001647 \\ -0.003552 & -0.002664 \end{bmatrix}$$

$$W_{hid}^n \approx \begin{bmatrix} 0.2011 & 0.4008 \\ 0.3018 & 0.1013 \end{bmatrix}$$

และ

$$B_{hidden}^n = B_{hidde} - \eta \frac{\partial E}{\partial B_{hidden}} = \begin{bmatrix} 0.1 \\ 0.2 \end{bmatrix} - 0.5 \begin{bmatrix} -0.000549 \\ -0.000888 \end{bmatrix}$$

$$B_{hidd}^n \approx \begin{bmatrix} 0.1003 \\ 0.2004 \end{bmatrix}$$

หมายเหตุ หลังจากการอัปเดต 1 รอบ เราจะได้ค่าอัปเดตซึ่งจะถูกนำไปใช้เป็นข้อมูลนำเข้าในการทำรอบที่ 2 การอัปเดตจะทำต่อเนื่องและจะหยุดเมื่อทำการเปรียบเทียบค่าอัปเดตในรอบปัจจุบันกับรอบก่อนหน้าและไม่มีการเปลี่ยนแปลงหรือมีการเปลี่ยนแปลงน้อยกว่าค่า threshold ที่ตั้งไว้ หรือเรียกอีกอย่างว่าการคำนวณเข้าสู่สภาวะเสถียร (steady state)

สรุป

การพยากรณ์อุปสงค์ในซัพพลายเชนยุคปัจจุบันไม่อาจอาศัยเพียงประสบการณ์หรือวิธีการเชิงสถิติแบบดั้งเดิมได้อีกต่อไป เพราะสภาพแวดล้อมทางธุรกิจเต็มไปด้วยความผันผวนและความไม่แน่นอน ทั้งจากการแข่งขันที่รุนแรงขึ้น พฤติกรรมผู้บริโภคที่เปลี่ยนแปลงรวดเร็ว การเติบโตของอีคอมเมิร์ซ ตลอดจนเหตุการณ์ภายนอกที่ส่งผลกระทบต่อการผลิต การขนส่ง และความต้องการของตลาดอย่างฉับพลัน ภายใต้บริบทเช่นนี้ องค์กรจำเป็นต้องใช้เครื่องมือที่สามารถเรียนรู้จากข้อมูลจำนวนมาก ค้นหารูปแบบที่ซับซ้อน และปรับตัวต่อความเปลี่ยนแปลงได้ดีกว่าวิธีเดิม การเรียนรู้ของเครื่องจึงกลายเป็นแนวทางสำคัญที่ช่วยยกระดับความแม่นยำของการพยากรณ์ และทำให้การตัดสินใจด้านการผลิต การจัดซื้อ การบริหารสินค้าคงคลัง และการกระจายสินค้ามีประสิทธิภาพมากยิ่งขึ้น

แนวคิดสำคัญของการใช้ ML ในการพยากรณ์อุปสงค์อยู่ที่ความสามารถในการนำข้อมูลจริงมาใช้สร้างแบบจำลองเพื่ออธิบายและคาดการณ์พฤติกรรมของอุปสงค์ในอนาคต ไม่ว่าจะเป็นข้อมูลยอดขายย้อนหลัง ราคา การส่งเสริมการขาย รายได้ของผู้บริโภค ฤดูกาล สภาพอากาศ หรือปัจจัยแวดล้อมอื่น ๆ ที่เกี่ยวข้อง วิธีการเหล่านี้มีข้อได้เปรียบตรงที่ไม่จำเป็นต้องจำกัดอยู่กับความสัมพันธ์แบบง่ายเสมอไป แต่สามารถเรียนรู้ความสัมพันธ์ที่ไม่เป็นเชิงเส้นและซับซ้อนกว่านั้นได้ จึงเหมาะกับโลกธุรกิจที่ข้อมูลมีมิติหลากหลายและเปลี่ยนแปลงอยู่ตลอดเวลา การทำความเข้าใจหลักการของแบบจำลองแต่ละประเภทจึงเป็นพื้นฐานสำคัญก่อนนำไปประยุกต์ใช้จริง

การถดถอยเชิงเส้นเป็นจุดเริ่มต้นที่สำคัญของการวิเคราะห์เชิงพยากรณ์ เพราะช่วยอธิบายความสัมพันธ์ระหว่างตัวแปรอิสระกับตัวแปรตามได้อย่างตรงไปตรงมา เช่น ความสัมพันธ์ระหว่างราคา โปรโมชัน และรายได้ครัวเรือนกับยอดขายสินค้า จุดเด่นของวิธีนี้คือความชัดเจนในการตีความผลลัพธ์ ผู้วิเคราะห์สามารถอธิบายได้ว่าปัจจัยใดส่งผลต่อยอดขายในทิศทางบวกหรือลบ และส่งผลมากน้อยเพียงใด จึงเหมาะสำหรับสถานการณ์ที่ต้องการทั้งการพยากรณ์และการอธิบายเหตุผลประกอบการตัดสินใจ อย่างไรก็ตาม การใช้การถดถอยเชิงเส้นให้ได้ผลดีจำเป็นต้องตรวจสอบสมมติฐานหลายประการ เช่น ความเป็นเชิงเส้นของความสัมพันธ์ ความเป็นอิสระของค่าคลาดเคลื่อน ความแปรปรวนคงที่ และการแจกแจงแบบปกติของค่าคลาดเคลื่อน หากข้อมูลไม่สอดคล้องกับเงื่อนไขเหล่านี้ ความแม่นยำและความน่าเชื่อถือของแบบจำลองก็อาจลดลงได้

เมื่อข้อมูลมีความซับซ้อนมากขึ้น แบบจำลองป่าสุ่มหรือ Random Forest จึงเข้ามามีบทบาทสำคัญในฐานะวิธีที่สามารถจัดการความสัมพันธ์ที่ไม่เป็นเชิงเส้นได้ดีกว่า แนวคิดของวิธีนี้คือการสร้างต้นไม้ตัดสินใจจำนวนมากจากชุดข้อมูลที่สุ่มแตกต่างกัน แล้วรวมผลลัพธ์จากต้นไม้เหล่านั้นเป็นค่าพยากรณ์สุดท้าย การรวมผลจากหลายต้นไม้ช่วยลดความผันผวนของการพยากรณ์และลด

ปัญหาการยึดติดกับข้อมูลเฉพาะชุดมากเกินไป ทำให้แบบจำลองมีเสถียรภาพมากขึ้น Random Forest จึงเหมาะกับข้อมูลจริงในงานซัพพลายเชน ซึ่งมักมีปัจจัยหลายด้านเข้ามาเกี่ยวข้องพร้อมกัน เช่น ยอดขายในอดีต ราคา โปรโมชัน ฤดูกาล สภาพอากาศ และพฤติกรรมของลูกค้า นอกจากนี้ใช้พยากรณ์ได้อย่างมีประสิทธิภาพแล้ว ยังช่วยให้มองเห็นว่าปัจจัยใดมีอิทธิพลต่ออุปสงค์มากที่สุดอีกด้วย

ในกรณีที่รูปแบบของข้อมูลมีความซับซ้อนสูงมาก โครงข่ายประสาทเทียมหรือ ANN เป็นอีกทางเลือกหนึ่งที่มีศักยภาพสูง แบบจำลองลักษณะนี้ได้รับแรงบันดาลใจจากการทำงานของสมองมนุษย์ โดยใช้ชั้นของหน่วยประมวลผลเชื่อมโยงกันผ่านค่าน้ำหนักที่ปรับเปลี่ยนได้ระหว่างกระบวนการเรียนรู้ จุดเด่นสำคัญคือความสามารถในการเรียนรู้ความสัมพันธ์ที่ซับซ้อนและซ่อนอยู่ในข้อมูลจำนวนมาก ซึ่งวิธีเชิงเส้นทั่วไปอาจไม่สามารถตรวจจับได้อย่างมีประสิทธิภาพ จึงเหมาะกับโจทย์พยากรณ์ที่มีตัวแปรจำนวนมากและมีปฏิสัมพันธ์ระหว่างกันหลายระดับ อย่างไรก็ตาม ANN ก็มีข้อจำกัดในด้านความซับซ้อนของการออกแบบแบบจำลอง ความต้องการข้อมูลจำนวนมาก และความยากในการตีความผลลัพธ์เมื่อเทียบกับวิธีที่เรียบง่ายกว่า ดังนั้น แม้จะเป็นเครื่องมือที่ทรงพลัง แต่ก็ไม่ได้หมายความว่าเหมาะสมกับทุกสถานการณ์เสมอไป

แบบฝึกหัด

จงตอบคำถามต่อไปนี้ ส่งงานกลุ่มที่ email: natapat.ar@ssru.ac.th

- กำหนดสมการถดถอยสำหรับการพยากรณ์ยอดขายสินค้า

$$Y = 247.757 - 47.048X_1 + 146.801X_2 + 0.042X_3$$

โดยที่

X_1 = ราคา (บาท)

X_2 = โปรโมชัน (1 = มีโปรโมชัน, 0 = ไม่มี)

X_3 = รายได้ครัวเรือน (บาท)

จงตอบคำถามต่อไปนี้

- 1.1 คำนวณยอดขายเมื่อ ราคา = 12 บาท โปรโมชัน = 0 รายได้ครัวเรือน = 32,000 บาท
- 1.2 คำนวณยอดขายเมื่อ ราคา = 12 บาท โปรโมชัน = 1 รายได้ครัวเรือน = 32,000 บาท
- 1.3 คำนวณ ผลกระทบเชิงส่วนเพิ่ม (Marginal Effect) ของโปรโมชันต่อยอดขาย ถ้าราคาเพิ่มขึ้น 2 บาท ยอดขายจะเปลี่ยนแปลงเท่าไร
- 1.4 อธิบายเชิงธุรกิจว่าผู้จัดการควรใช้ข้อมูลจากโมเดลนี้เพื่อกำหนดกลยุทธ์ราคาอย่างไร

- สมมติว่าแบบจำลอง Random Forest ประกอบด้วย 5 Decision Trees ที่ให้ค่าพยากรณ์ยอดขายดังนี้

Tree	1	2	3	4	5
Prediction	128	134	130	140	138

จงตอบคำถามต่อไปนี้

- 2.1 ค่าพยากรณ์สุดท้ายของ Random Forest
- 2.2 คำนวณค่า Variance ของ Prediction จากต้นไม้ทั้ง 5 ต้น
- 2.3 อธิบายว่าทำไมการรวมผลลัพธ์จากหลายต้นไม้จึงช่วยลด Variance ของโมเดล
- 2.4 ถ้ามีต้นไม้เพิ่มเป็น 100 ต้น โดยค่าเฉลี่ยไม่เปลี่ยนแปลง โมเดลจะมีข้อดีอย่างไร

3. กำหนดโครงข่ายประสาทเทียมดังนี้

Input

$$x_1 = 5, x_2 = 2$$

Hidden layer มี 2 Neurons

น้ำหนัก

$$W_{hidden} = \begin{bmatrix} 0.3 & 0.2 \\ 0.4 & 0.1 \end{bmatrix}$$

Bias

$$B_{hidden} = \begin{bmatrix} 0.1 \\ 0.2 \end{bmatrix}$$

Output layer

$$W_{output} = \begin{bmatrix} 0.6 & 0.5 \end{bmatrix}$$
$$B_{output} = 0.1$$

ให้ใช้ Sigmoid Activation

จงคำนวณ

1. ค่า Z_1 ของ Hidden Layer ในรูป Matrix
2. ค่า $A(Z_1)$ หลัง Sigmoid
3. ค่า Z_2 ของ Output Layer
4. ค่า Output สุดท้าย

เอกสารอ้างอิง

- Areerakulkan, N., & Pongpech, W. (2021). A Dempster-Shafer big data readiness assessment model. In *Proceedings of the 23rd International Conference on Enterprise Information Systems (ICEIS 2021)* (Vol. 2, pp. 581-585). SCITEPRESS. <https://doi.org/10.5220/0010506205810585>.
- Bandara, K., Bergmeir, C., & Smyl, S. (2020). Forecasting across time series databases using recurrent neural networks on groups of similar series: A clustering approach. *Expert Systems with Applications*, 140, 112896. <https://doi.org/10.1016/j.eswa.2019.112896>.
- Baryannis, G., Dani, S., & Antoniou, G. (2019). Predicting supply chain risks using machine learning: The trade-off between performance and interpretability. *Future Generation Computer Systems*, 101, 993–1004. <https://doi.org/10.1016/j.future.2019.07.059>
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5-32. <https://doi.org/10.1023/A:1010933404324>.
- Carbonneau, R., Laframboise, K., & Vahidov, R. (2008). Application of machine learning techniques for supply chain demand forecasting. *European Journal of Operational Research*, 184(3), 1140-1154. <https://doi.org/10.1016/j.ejor.2006.12.004>.
- Choi, T. M., Wallace, S. W., & Wang, Y. (2018). Big data analytics in operations management. *Production and Operations Management*, 27(10), 1868-1883. <https://doi.org/10.1111/poms.12838>.
- Genuer, R., Poggi, J. M., Tuleau-Malot, C., & Villa-Vialaneix, N. (2015). Random forests for big data. *arXiv preprint arXiv:1511.08327*. <https://arxiv.org/abs/1511.08327>.
- Molnar, C. (2022) Interpretable Machine Learning. A Guide for Making Black Box Models Explainable, Self-Published. <https://christophm.github.io/interpretable-ml-book/>
- Islam, S., & Amin, S. H. (2020). Prediction of probable backorder scenarios in the supply chain using distributed random forest and gradient boosting. *Journal of Big Data*, 7, 65. <https://doi.org/10.1186/s40537-020-00345-2>.

เอกสารอ้างอิง (ต่อ)

- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An introduction to statistical learning: With applications in R* (2nd ed.). Springer.
- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2018). Statistical and machine learning forecasting methods: Concerns and ways forward. *PLOS ONE*, 13(3), e0194889. <https://doi.org/10.1371/journal.pone.0194889>.
- Mediavilla, M. A., Escudero-Santana, A., & Nieto-Morote, A. (2022). Review and analysis of artificial intelligence methods for demand forecasting in supply chain. *Procedia Computer Science*, 200, 112-121. <https://doi.org/10.1016/j.procir.2022.05.119>.
- Mitra, A., Jain, A., Kishore, A. et al. A Comparative Study of Demand Forecasting Models for a Multi-Channel Retail Company: A Novel Hybrid Machine Learning Approach. *Oper. Res. Forum* 3, 58 (2022). <https://doi.org/10.1007/s43069-022-00166-4>.
- Scornet, E., Biau, G., & Vert, J. P. (2015). Consistency of random forests. *The Annals of Statistics*, 43(4), 1716-1741. <https://doi.org/10.1214/15-AOS1321>.
- Schonlau, M., & Zou, R. Y. (2020). The random forest algorithm for statistical learning. *The Stata Journal*, 20(1), 3-29. <https://doi.org/10.1177/1536867X20909688>.
- Van Steenbergen, R. M., Mes, M. R. K., & Van Harten, A. (2020). Forecasting demand profiles of new products. *European Journal of Operational Research*, 287(2), 555-568. <https://doi.org/10.1016/j.ejor.2020.04.014>.
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77-84. <https://doi.org/10.1111/jbl.12010>
- Zhang, G., Patuwo, B. E., & Hu, M. Y. (1998). Forecasting with artificial neural networks: The state of the art. *International Journal of Forecasting*, 14(1), 35-62. [https://doi.org/10.1016/S0169-2070\(97\)00044-7](https://doi.org/10.1016/S0169-2070(97)00044-7).



ส่วนที่ 4

การวิเคราะห์เพื่อการตัดสินใจ

Part 4: Prescriptive Analytics

“ข้อมูลสร้างความเข้าใจ
การคาดการณ์สร้างความพร้อม
และการตัดสินใจสร้างความได้เปรียบ”

บทที่ 7

การจัดการและวิเคราะห์ข้อมูลโลจิสติกส์

Logistics Management and Logistics Data Analytics

หัวข้อ

- 7.1 การจัดการโลจิสติกส์
- 7.2 กิจกรรมหลักในการจัดการโลจิสติกส์
- 7.3 รูปแบบการขนส่ง
- 7.4 การออกแบบเครือข่ายโลจิสติกส์
- 7.5 การออกแบบเครือข่ายโลจิสติกส์โดยใช้แบบจำลองเชิงเส้น (Linear Programming, LP)

การจัดการโลจิสติกส์ (logistics management) คือกระบวนการวางแผน ดำเนินการ และควบคุมการเคลื่อนย้ายและจัดเก็บสินค้า วัตถุดิบ ข้อมูล และทรัพยากรต่าง ๆ จากจุดเริ่มต้นถึงปลายทาง โดยมุ่งเน้นให้เกิดความคุ้มค่าในเชิงต้นทุนและเวลา รวมถึงการตอบสนองต่อความต้องการของลูกค้าได้อย่างมีประสิทธิภาพ การจัดการโลจิสติกส์ครอบคลุมการจัดการสินค้าคงคลัง การขนส่ง การจัดเก็บสินค้า การดำเนินการคลังสินค้า และการจัดการข้อมูลและระบบสารสนเทศที่เกี่ยวข้อง การจัดการโลจิสติกส์มีความสำคัญในหลายด้าน ได้แก่ (1) เพิ่มประสิทธิภาพการดำเนินงาน การจัดการโลจิสติกส์ที่ดีช่วยให้กระบวนการต่าง ๆ ในซัพพลายเชนทำงานได้อย่างราบรื่น ลดความล่าช้าและความสูญเสีย ทำให้สามารถตอบสนองต่อความต้องการของตลาดได้อย่างรวดเร็วและมีประสิทธิภาพ (2) ช่วยลดต้นทุนในหลายด้าน เช่น การขนส่ง การจัดเก็บ การดำเนินการคลังสินค้า และการบริหารสินค้าคงคลัง การลดต้นทุนเหล่านี้มีผลโดยตรงต่อความสามารถในการแข่งขันขององค์กร (3) ปรับปรุงการบริการลูกค้าช่วยให้บริษัทสามารถส่งมอบสินค้าและบริการได้ตรงเวลา และในสภาพที่ดี ทำให้ลูกค้ามีความพึงพอใจและสร้างความเชื่อมั่นในผลิตภัณฑ์และบริการของบริษัท ซึ่งมีผลต่อความภักดีของลูกค้าและชื่อเสียงของบริษัท (4) การจัดการความเสี่ยง การจัดการโลจิสติกส์มีบทบาทสำคัญในการบริหารความเสี่ยงที่เกี่ยวข้องกับการเคลื่อนย้ายสินค้า เช่น การจัดการกับสถานการณ์ที่อาจเกิดขึ้นได้ เช่น ภัยธรรมชาติ การเปลี่ยนแปลงของกฎระเบียบ หรือความล่าช้าในกระบวนการขนส่ง (5) สนับสนุนการตัดสินใจที่มีข้อมูล ข้อมูลที่เกิดจากการจัดการโลจิสติกส์สามารถ

นำมาใช้ในการวางแผนและตัดสินใจในด้านต่าง ๆ ของธุรกิจ เช่น การพยากรณ์ความต้องการ การวางแผนการผลิต และการบริหารสินค้าคงคลัง ซึ่งช่วยให้บริษัทสามารถตัดสินใจได้อย่างมีประสิทธิภาพและทันเวลา (6) สนับสนุนการขยายธุรกิจ การจัดการโลจิสติกส์ที่มีประสิทธิภาพช่วยให้บริษัทสามารถขยายธุรกิจไปยังตลาดใหม่ ๆ ได้อย่างรวดเร็วและมั่นใจ โดยลดอุปสรรคที่เกี่ยวข้องกับการขนส่งและการจัดการสินค้าข้ามประเทศ

สำหรับการวิเคราะห์ข้อมูลโลจิสติกส์ (logistics data analytics) เป็นกระบวนการที่นำข้อมูลที่ได้จากกิจกรรมโลจิสติกส์ต่าง ๆ มาวิเคราะห์เพื่อสร้างความเข้าใจและตัดสินใจที่มีประสิทธิภาพมากยิ่งขึ้นในกระบวนการจัดการโลจิสติกส์ การวิเคราะห์ข้อมูลโลจิสติกส์เกี่ยวข้องกับการรวบรวม ประมวลผล และตีความข้อมูลเพื่อปรับปรุงการดำเนินงานในซัพพลายเชน ซึ่งเป็นส่วนสำคัญในการเพิ่มประสิทธิภาพ ลดต้นทุน และปรับปรุงการบริการลูกค้า ดังนั้นการจัดการและการวิเคราะห์ข้อมูลโลจิสติกส์ไม่ใช่เพียงแค่การเคลื่อนย้ายสินค้า แต่ยังเป็นส่วนสำคัญในการสร้างมูลค่าให้กับองค์กร และช่วยให้องค์กรสามารถตอบสนองต่อความเปลี่ยนแปลงในตลาดได้อย่างรวดเร็วและมีประสิทธิภาพ การจัดการและการวิเคราะห์ข้อมูลโลจิสติกส์ที่ดีและถูกต้องแม่นยำจึงเป็นปัจจัยสำคัญที่ส่งเสริมความสำเร็จและการเติบโตของธุรกิจในระยะยาว (Bowersox et. al., 2019; Chopra & Meindl, 2019)

7.1 การจัดการโลจิสติกส์

การจัดการโลจิสติกส์ถูกพิจารณาว่าเป็นส่วนหนึ่งของกระบวนการในซัพพลายเชนที่เกี่ยวข้องกับการเคลื่อนย้ายและกระจายสินค้าบริการ (Rushton et al., 2022) ตามที่สถาบัน Chartered Institute of Purchasing & Supply (CIPS) ได้ให้คำจำกัดความไว้ว่า

"การจัดการโลจิสติกส์เป็นส่วนประกอบของการจัดการซัพพลายเชน ใช้เพื่อตอบสนองความต้องการของลูกค้าผ่านการวางแผน การควบคุม และการดำเนินการเคลื่อนย้ายสินค้าอย่างมีประสิทธิภาพ"

การจัดการโลจิสติกส์มีบทบาทสำคัญในการจัดการซัพพลายเชนอย่างมีประสิทธิภาพ (Mentzer et al., 2004) เนื่องจากการจัดการโลจิสติกส์ที่ดีช่วยให้การไหลของวัสดุภายในซัพพลายเชนเป็นไปอย่างราบรื่นและทันเวลาเพื่อตอบสนองความต้องการของลูกค้า อุตสาหกรรมโลจิสติกส์ได้กลายเป็นกระดูกสันหลังของเศรษฐกิจ

ไม่เพียงแต่การจัดส่งสินค้าจากศูนย์กลางและท่าเรือเท่านั้น แต่ยังช่วยสนับสนุนอุตสาหกรรมอื่น ๆ เช่น การผลิต การค้าปลีก การก่อสร้าง และบริการ

อย่างไรก็ตามในสภาพแวดล้อมทางธุรกิจที่มีการแข่งขันสูงในปัจจุบัน ผู้ค้าปลีกหลายราย กำลังเปลี่ยนแปลงการดำเนินงานจากธุรกิจแบบร้านค้าปลีกดั้งเดิมไปเป็นการช้อปปิ้งออนไลน์ (Ivanov & Dolgui, 2021) หนึ่งในปัจจัยสำคัญที่สนับสนุนความสำเร็จของการเปลี่ยนแปลงนี้คือการดำเนินงานด้านโลจิสติกส์ที่รวดเร็วและมีประสิทธิภาพ ตัวอย่างเช่น เหตุผลที่ Amazon เหนือกว่า คู่แข่งหลายรายเกิดจากระบบโลจิสติกส์ที่มีประสิทธิภาพซึ่งสามารถจัดส่งสินค้าถึงลูกค้าในวันถัดไป ในวันเดียวกัน หรือแม้กระทั่งเร็วกว่านั้น

หากการจัดการซัพพลายเชนมีผลกระทบอย่างมีนัยสำคัญต่อธุรกิจที่ช่วยให้ชนะคู่แข่งของตน การจัดการโลจิสติกส์ก็มีบทบาทสำคัญด้านประสิทธิภาพของบริษัทในแง่ของต้นทุน คุณภาพ ความเร็ว ความยืดหยุ่น และความน่าเชื่อถือ (Harrison et al., 2019) ตารางที่ 7.1 แสดงตัวอย่างบางส่วนของการจัดการโลจิสติกส์ที่เป็นส่วนเสริมให้บรรลุวัตถุประสงค์ด้านประสิทธิภาพสำหรับการจัดการซัพพลายเชนในภาพรวม

ตารางที่ 7.1 บทบาทการจัดการโลจิสติกส์ที่ช่วยให้บรรลุวัตถุประสงค์ด้านประสิทธิภาพ

วัตถุประสงค์	บทบาทของการจัดการโลจิสติกส์
ต้นทุน	การเลือกโหมดการขนส่งที่เหมาะสม และการจัดการยานพาหนะอย่างมีประสิทธิภาพสามารถช่วยลดต้นทุนได้
คุณภาพ	การจัดส่งที่รวดเร็ว แม่นยำ และตรงเวลาหมายถึงคุณภาพสูงในสายตาของลูกค้าที่รอคอยสินค้ามาถึง
ความเร็ว	การเลือกตำแหน่งที่ตั้งของศูนย์กระจายสินค้า (DC) และการออกแบบเครือข่ายอย่างมีประสิทธิภาพสามารถช่วยให้การดำเนินการคำสั่งซื้อเป็นไปอย่างรวดเร็ว
ความยืดหยุ่น	การดำเนินงานด้านโลจิสติกส์อย่างมีประสิทธิภาพสามารถรองรับคำขอของลูกค้าที่แตกต่างกันได้ในแง่ของปริมาณสินค้า ประเภทสินค้า และการเปลี่ยนแปลงของช่วงเวลาการจัดส่ง
ความน่าเชื่อถือ	การออกแบบเครือข่ายศูนย์กระจายสินค้า (DC) และการจัดการยานพาหนะอย่างมีประสิทธิภาพสามารถช่วยเสริมสร้างการจัดส่งที่น่าเชื่อถือ ปลอดภัย และแม่นยำให้กับลูกค้า

ที่มา: จัดทำโดยผู้แต่ง

7.2 กิจกรรมหลักในการจัดการโลจิสติกส์

ภายในซัพพลายเชนทั่วไป การจัดการโลจิสติกส์สามารถแบ่งออกเป็นสองประเภทใหญ่ๆ ได้แก่ โลจิสติกส์ขาเข้า (inbound logistics) และโลจิสติกส์ขาออก (outbound logistics) โดยที่โลจิสติกส์ขาเข้าจะจัดการกับการรับวัตถุดิบ สินค้า และชิ้นส่วนจากซัพพลายเออร์ต้นน้ำไปยังธุรกิจหลัก ส่วนโลจิสติกส์ขาออกครอบคลุมการส่งมอบสินค้าสำเร็จรูปหรือกึ่งสำเร็จรูปไปยังลูกค้าปลายทาง ซึ่งรวมถึงศูนย์กระจายสินค้า (DCs) ผู้ค้าส่ง ผู้ค้าปลีก และผู้บริโภคสุดท้าย (end consumers)

กิจกรรมหลักที่ทั้งการจัดการโลจิสติกส์ขาเข้าและขาออกเกี่ยวข้อง ได้แก่ การจัดการคลังสินค้า การจัดเก็บ การจัดการวัสดุ การบรรจุภัณฑ์ การขนส่ง รวมถึงข้อมูลและความปลอดภัย (Grant et al., 2022) โดยมีรายละเอียดโดยสรุปดังนี้

1. การบริหารจัดการการขนส่ง (Transportation Management) เป็นการจัดการกระบวนการขนส่งสินค้าโดยใช้วิธีที่เหมาะสมที่สุด ทั้งในด้านของประเภทการขนส่ง (ทางบก ทางน้ำ ทางอากาศ) การจัดการเส้นทาง และการประสานงานกับผู้ขนส่ง เพื่อให้สินค้าถึงปลายทางได้ตามกำหนดเวลาและด้วยต้นทุนที่ต่ำที่สุด

2. การบริหารจัดการคลังสินค้า (Warehouse Management) เป็นการจัดเก็บสินค้าในคลังอย่างมีประสิทธิภาพ ทั้งในแง่ของการจัดพื้นที่ การควบคุมปริมาณสต็อก การรับและจัดส่งสินค้า รวมถึงการจัดเก็บสินค้าในลักษณะที่เหมาะสมกับประเภทสินค้า เพื่อให้สามารถดึงสินค้าส่งมอบได้อย่างรวดเร็วและแม่นยำ

3. การบริหารจัดการสินค้าคงคลัง (Inventory Management) เป็นการควบคุมและจัดการปริมาณสินค้าคงคลังให้อยู่ในระดับที่เหมาะสมเพื่อลดต้นทุนในขณะที่ยังสามารถตอบสนองความต้องการของลูกค้าได้อย่างทันเวลา ซึ่งรวมถึงการพยากรณ์ความต้องการ การเติมสินค้า และการควบคุมการจัดเก็บสินค้า

4. การจัดการข้อมูลและการสื่อสาร (Information and Communication Management) โลจิสติกส์ต้องการข้อมูลที่แม่นยำและทันเวลาในการติดตามการเคลื่อนย้ายสินค้า การสื่อสารที่มีประสิทธิภาพระหว่างผู้ผลิต ผู้จัดจำหน่าย และลูกค้า รวมถึงการใช้เทคโนโลยีสารสนเทศเช่น ระบบจัดการการขนส่ง (TMS) และระบบจัดการคลังสินค้า (WMS) เพื่อสนับสนุนการทำงาน (Ivanov & Dolgui, 2020)

5. การจัดการกระบวนการคืนสินค้า (Reverse Logistics) เป็นการจัดการกระบวนการนำสินค้ากลับมาจากลูกค้า ซึ่งอาจเป็นการคืนสินค้าเนื่องจากปัญหาด้านคุณภาพ การรับคืนจากการขาย หรือการรีไซเคิลสินค้า โดยต้องมีการวางแผนให้กระบวนการนี้มีความคล่องตัวและประหยัด (Govindan & Soleimani, 2017)

6. การบริการลูกค้า (Customer Service) การบริการลูกค้าในโลจิสติกส์เน้นที่การตอบสนองความต้องการของลูกค้าในเรื่องของการจัดส่งสินค้า ความถูกต้องของคำสั่งซื้อ และการแก้ไขปัญหาที่เกิดขึ้น เพื่อรักษาความพึงพอใจของลูกค้าในระยะยาว

นอกจากกิจกรรมหลักดังที่กล่าวมาแล้ว การออกแบบเครือข่ายโลจิสติกส์ก็ถือว่าเป็นกิจกรรมสำคัญในการจัดการโลจิสติกส์ในองค์กรขนาดใหญ่ที่มีโครงสร้างของซัพพลายเชนที่ซับซ้อนเกี่ยวกับซัพพลายเออร์ต้นน้ำหลายรายรวมไปถึงศูนย์กระจายสินค้าปลายน้ำ และผู้ค้าปลีก เป็นต้น โดยที่การออกแบบเครือข่ายโลจิสติกส์สามารถแบ่งออกเป็นสองส่วนหลัก ๆ ได้ดังนี้

1. การตัดสินใจเกี่ยวกับตำแหน่งที่ตั้ง คือ การเลือกตำแหน่งของโหนดในซัพพลายเชน เช่น ฐานการผลิต ศูนย์กระจายสินค้า (DCs) และเครือข่ายค้าปลีก โดยพิจารณาตำแหน่งของลูกค้าเป็นหลัก การตัดสินใจเกี่ยวกับตำแหน่งของซัพพลายเออร์โดยพิจารณาจากเวลานำส่ง ต้นทุน และเกณฑ์อื่นๆ ตัวอย่างเช่น การจัดส่งแบบทันเวลาพอดี (Just-in-time - JIT) เป็นการปฏิบัติด้านโลจิสติกส์ทั่วไปในอุตสาหกรรมยานยนต์ ซึ่งต้องการให้ซัพพลายเออร์ (หรือคลังสินค้าของพวกเขา) ตั้งอยู่ในระยะใกล้เคียงกับบริษัท

2. การจัดเส้นทาง คือ การเลือกและวางแผนเส้นทางจัดส่งด้วยวัตถุประสงค์ต่าง ๆ เช่น การลดระยะทางการเดินทางทั้งหมด ต้นทุน เวลา เป็นต้น ตัวอย่างเช่น บริษัทขนส่งพัสดุแห่งหนึ่งต้องการส่งพัสดุให้กับลูกค้า 10 จุดหมายภายในเมืองเดียวกัน การวางแผนเส้นทางจะต้องพิจารณาปัจจัยที่เกี่ยวข้องได้แก่ ระยะทางระหว่างจุดหมายแต่ละแห่ง (เส้นทางจัดส่งควรเป็นเส้นทางที่สั้นที่สุดระหว่างทุกจุดหมาย) ลำดับการจัดส่ง (ควรกำหนดลำดับให้รถขนส่งวิ่งไปในทิศทางที่ทำให้ระยะทางขนส่งรวมทั้งหมดสั้นที่สุด) และการใช้ยานพาหนะอย่างมีประสิทธิภาพ (ควรใช้งานรถขนส่งอย่างคุ้มค่าโดยการกำหนดจำนวนจุดหมายที่ต้องจัดส่งให้เหมาะสม)

7.3 รูปแบบการขนส่ง

ในการจัดการโลจิสติกส์จะมีรูปแบบการขนส่งอยู่ 5 รูปแบบ ได้แก่ การขนส่ง ทางบก ทางราง ทางน้ำ ทางอากาศ และ ทางท่อ แต่ละโหมดการขนส่งมีวัตถุประสงค์ในการใช้งานที่แตกต่างกันไป รวมถึงข้อดีและข้อเสียตามรายละเอียดในตารางที่ 7.2 การเลือกโหมดการขนส่งขึ้นอยู่กับหลายปัจจัย เช่น ต้นทุน ความพร้อมของบริการ ความเร็ว ความยืดหยุ่น ความเสี่ยง และการพิจารณาด้านความปลอดภัย (Rodrigue, 2020) ตัวอย่างเช่น การขนส่งทางถนนด้วยรถบรรทุกและรถตู้สามารถให้ความยืดหยุ่นสูงสุดในการจัดส่งจากต้นทางถึงปลายทาง และมีความสะดวกในการเข้าถึงสูงเมื่อเทียบกับโหมดการขนส่งอื่น ๆ อย่างไรก็ตาม สำหรับระยะทางที่ไกลขึ้น โหมดนี้อาจจะช้ากว่าและมีต้นทุนค่อนข้างสูง เช่น การขนส่งสินค้าจากจีนไปยังสหราชอาณาจักร ดังนั้น ควรพิจารณาโหมดการขนส่งทางเลือก เช่น การขนส่งทางทะเลหรือเรือข้ามทางราง เป็นต้น รายละเอียดเพิ่มเติมของโหมดการขนส่งแต่ละโหมดอธิบายดังต่อไปนี้ (Novack et al., 2019)

1. การขนส่งทางถนน (Road Transportation) การขนส่งทางถนนเป็นวิธีที่ใช้บ่อยที่สุด โดยเฉพาะในกรณีของการขนส่งสินค้าภายในประเทศ การขนส่งทางถนนมีความยืดหยุ่นสูง สามารถส่งสินค้าไปถึงสถานที่ต่าง ๆ ได้โดยตรง ทำให้เหมาะสำหรับการขนส่งระยะสั้นและการขนส่งสินค้าที่มีน้ำหนักไม่มาก

2. การขนส่งทางราง (Rail Transportation) การขนส่งทางรางมีประสิทธิภาพสูงในการขนส่งสินค้าปริมาณมากในระยะทางไกล เช่น สินค้าโภคภัณฑ์ สินค้าเกษตร และวัสดุก่อสร้าง การขนส่งทางรางมีต้นทุนต่อหน่วยที่ต่ำ แต่มีข้อจำกัดเรื่องความยืดหยุ่นในการเข้าถึงสถานที่ต่าง ๆ

3. การขนส่งทางน้ำ (Water Transportation) การขนส่งทางน้ำเป็นวิธีที่มีประสิทธิภาพสูงในการขนส่งสินค้าปริมาณมากระหว่างประเทศ เช่น สินค้าโภคภัณฑ์และสินค้าขนาดใหญ่ โดยเฉพาะการขนส่งระหว่างทวีป อย่างไรก็ตาม การขนส่งทางน้ำมีความเร็วต่ำและขึ้นอยู่กับสภาพอากาศ

4. การขนส่งทางอากาศ (Air Transportation) การขนส่งทางอากาศมีความเร็วสูงที่สุด จึงเหมาะสำหรับการขนส่งสินค้าที่มีมูลค่าสูง สินค้าที่มีความเร่งด่วน หรือสินค้าที่เน่าเสียได้เร็ว เช่น อิเล็กทรอนิกส์และยา แต่การขนส่งทางอากาศมีต้นทุนสูงที่สุดเมื่อเทียบกับวิธีการอื่น ๆ

5. การขนส่งทางท่อ (Pipeline Transportation) การขนส่งทางท่อใช้สำหรับการขนส่งของเหลวและก๊าซ เช่น น้ำมันและก๊าซธรรมชาติ แม้จะมีข้อจำกัดในประเภทของสินค้าที่ขนส่งได้ แต่เป็นวิธีที่มีต้นทุนต่ำและมีประสิทธิภาพในระยะทางไกล

ดังที่กล่าวมาแล้วข้างต้น การเลือกโหมดการขนส่งที่เหมาะสมขึ้นอยู่กับหลายปัจจัย เช่น ลักษณะของสินค้า ระยะทาง ต้นทุน เวลาที่ต้องใช้ในการขนส่ง แสดงดังตารางที่ 7.2 และความต้องการของลูกค้า โดยการวิเคราะห์และเลือกโหมดการขนส่งที่เหมาะสมจะช่วยเพิ่มประสิทธิภาพในซัพพลายเชนและลดต้นทุนโลจิสติกส์โดยรวม

ตารางที่ 7.2 ตารางเปรียบเทียบปัจจัยการเลือกโหมดการขนส่ง

โหมดการขนส่ง	ข้อดี	ข้อเสีย	ต้นทุน (ต่อหน่วย)	เวลา	ความยืดหยุ่น	ประเภทสินค้าที่เหมาะสม
ทางถนน (Road)	ยืดหยุ่นสูงส่งได้ถึงที่หมายโดยตรง	จำกัดด้วยระยะทางและปริมาณสินค้าที่ขนส่งได้	ปานกลาง	ปานกลาง	สูง	สินค้าทั่วไป สินค้าอุปโภคบริโภค
ทางราง (Rail)	ขนส่งได้ปริมาณมากในระยะไกล	ยืดหยุ่นต่ำ ต้องใช้ร่วมกับโหมดอื่น ๆ	ต่ำ	ช้า	ต่ำ	สินค้าหนัก วัตถุดิบ โภคภัณฑ์
ทางน้ำ (Water)	ประหยัดต้นทุนสำหรับปริมาณสินค้าขนาดใหญ่	ความเร็วต่ำ ขึ้นอยู่กับสภาพอากาศ	ต่ำมาก	ช้ามาก	ต่ำ	สินค้าปริมาณมาก สินค้าขนาดใหญ่

ตารางที่ 7.2 (ต่อ)

โหมดการขนส่ง	ข้อดี	ข้อเสีย	ต้นทุน (ต่อหน่วย)	เวลา	ความยืดหยุ่น	ประเภทสินค้าที่เหมาะสม
ทางอากาศ (Air)	เร็วที่สุด เหมาะสำหรับสินค้าที่ต้องการความรวดเร็ว	ต้นทุนสูงมาก ข้อจำกัดในปริมาณและน้ำหนัก	สูงมาก	เร็วที่สุด	ปานกลาง	สินค้ามูลค่าสูง สินค้าเร่งด่วน ยา อิเล็กทรอนิกส์

ที่มา: จัดทำโดยผู้แต่ง

7.4 การออกแบบเครือข่ายโลจิสติกส์

การออกแบบเครือข่ายโลจิสติกส์เป็นหัวข้อที่กว้างขวางและสำคัญสำหรับการจัดการโลจิสติกส์และซัพพลายเชนอย่างมีประสิทธิภาพ (Simchi-Levi et al., 2021) เครือข่ายโลจิสติกส์ที่ได้รับการปรับปรุงสามารถลดต้นทุน ปรับปรุงการไหลของวัสดุ และเสริมสร้างการบริการแก่ลูกค้า การวิเคราะห์การตัดสินใจด้านการออกแบบเครือข่ายโลจิสติกส์ต่าง ๆ ได้แก่ (1) บทบาทของสถานที่ (facility role) บทบาทและกระบวนการที่แต่ละสถานที่ควรมีใน เครือข่ายคืออะไร? (2) ที่ตั้งของสถานที่ (facility location) ควรสร้างสถานที่ที่ไหน? (3) จำนวนสถานที่ (number of facilities) ควรสร้างสถานที่กี่แห่ง? (4) การจัดสรรความจุ/สมรรถนะ (capacity allocation) ควรจัดสรรความจุ/สมรรถนะให้แต่ละสถานที่เท่าไร? (5) การจัดสรรตลาด (market allocation) สถานที่แต่ละแห่งควรให้บริการตลาดใดบ้าง? (6) การเลือกเส้นทาง (route choice) ควรเลือกเส้นทางจัดส่งใดในวิธีที่ได้รับการปรับปรุง?

7.4.1 การวิเคราะห์ตัดสินใจเกี่ยวกับสถานที่ตั้ง

ในการออกแบบเครือข่ายโลจิสติกส์ การเลือกสถานที่อาจเป็นหนึ่งในการตัดสินใจที่สำคัญที่สุดที่ผู้จัดการด้านโลจิสติกส์และซัพพลายเชนต้องเผชิญ ตัวอย่างเช่น เมื่อเลือกซัพพลายเออร์ ปัจจัยสำคัญที่ต้องพิจารณาคือสถานที่ตั้งของพวกเขาและ หากจำเป็นซัพพลายเออร์เต็มใจที่จะสร้างโรงงานหรือคลังสินค้าใกล้กับที่ตั้งของคุณหรือไม่ เช่น เมื่อ Jaguar Land Rover (JLR) สร้างโรงงานแรกในประเทศจีน ซัพพลายเออร์สำคัญหลายรายก็ย้ายมาตั้งโรงงานใกล้ฐานการผลิตของ JLR ด้วย

นอกเหนือจากการเลือกซัพพลายเออร์ในต้นน้ำ การตัดสินใจเกี่ยวกับสถานที่ยังรวมถึงการตัดสินใจในปลายน้ำจากการเลือกสถานที่ผลิต ศูนย์กระจายสินค้า (DCs) คลังสินค้า จนถึงร้านค้าปลีกองค์กรมักต้องพิจารณาว่าควรสร้างสถานที่ใหม่ที่ไหนหากความต้องการเพิ่มขึ้น หรือปิดสถานที่ในตำแหน่งหนึ่งและย้ายไปยังอีกตำแหน่งหากความต้องการเปลี่ยนแปลง เห็นได้ชัดว่านี่เป็นการตัดสินใจเชิงกลยุทธ์เกี่ยวกับสถานที่ที่อาจมีผลกระทบระยะยาวต่อประสิทธิภาพของบริษัทที่มีปัจจัยหลายอย่างส่งผลต่อการตัดสินใจเกี่ยวกับสถานที่ เช่น ต้นทุน ระดับความต้องการในภูมิภาค และกฎระเบียบด้านล่างนี้คือปัจจัยสำคัญบางประการที่ควรพิจารณา (Farahani et al., 2010; Singh, et. al., 2018)

1. **ต้นทุน (Cost)** ได้แก่ ค่าที่ดินและค่าเช่า ต้นทุนแรงงาน และต้นทุนการขนส่ง
2. **ระดับความต้องการในภูมิภาค (Regional Demand Level)** ได้แก่ ความต้องการของลูกค้าในพื้นที่ และแนวโน้มการเติบโตของตลาดในพื้นที่ด้วย
3. **กฎระเบียบ (Regulations)** ได้แก่ ข้อบังคับและกฎหมายท้องถิ่น และนโยบายทางภาษีและการสนับสนุนจากรัฐบาล
4. **โครงสร้างพื้นฐาน (Infrastructure)** ได้แก่ ความพร้อมของโครงสร้างพื้นฐาน เช่น ถนน ท่าเรือ สนามบิน และการเข้าถึงแหล่งพลังงานและทรัพยากร
5. **ความใกล้ชิดกับซัพพลายเออร์และลูกค้า (Proximity to Suppliers and Customers)** จะช่วยลดเวลาการขนส่งและต้นทุนการขนส่ง และช่วยเพิ่มความยืดหยุ่นในการตอบสนองความต้องการของตลาด
6. **คุณภาพชีวิตและสิ่งแวดล้อม (Quality of Life and Environment)** เช่น สภาพแวดล้อมการทำงานที่ดีสำหรับพนักงาน ความยั่งยืนและผลกระทบต่อสิ่งแวดล้อม
7. **ระดับการบริการ (Service Level)** เกี่ยวข้องกับความใกล้ชิดกับตลาดและระดับความต้องการ เพื่อให้บรรลุระดับการบริการที่สูง บริษัทต้องพิจารณาไม่เพียงแต่สถานที่ที่จะ

สร้าง แต่ยักรวมถึงจำนวนและขนาดของสถานที่ที่จะสร้าง รวมถึงโหมดการขนส่งที่เหมาะสม

8. **ปัจจัยความเสี่ยง (Risk Factor)** อาจเกี่ยวข้องกับความปลอดภัยทางเมือง วัฒนธรรม เศรษฐกิจ และสิ่งแวดล้อมที่กว้างขึ้น รวมถึงความเสี่ยงจากการแข่งขัน ความผันแปรของความต้องการ และการสูญเสียลูกค้าที่อาจเกิดขึ้นเนื่องจาก การขาดแคลนสินค้า ความล่าช้า และกระบวนการคืนสินค้า (Ivanov, 2021)

การพิจารณาปัจจัยเหล่านี้จะช่วยให้ผู้จัดการด้านโลจิสติกส์และซัพพลายเชนสามารถตัดสินใจเกี่ยวกับสถานที่ที่ดีที่สุดสำหรับองค์กรของตนได้

7.4.2 วิธีศูนย์กลางความโน้มถ่วง (Centre of Gravity Method)

วิธีศูนย์กลางความโน้มถ่วงเป็นวิธีการทางคณิตศาสตร์ที่ใช้ในการกำหนดตำแหน่งที่ตั้งที่เหมาะสมที่สุดสำหรับศูนย์กระจายสินค้า (DC) หรือสถานที่อื่น ๆ ในซัพพลายเชน โดยพิจารณาจากตำแหน่งและปริมาณของศูนย์ที่มีอยู่และลูกค้าในเครือข่าย

ขั้นตอนของวิธีศูนย์กลางความโน้มถ่วง

1. **กำหนดพิกัดและปริมาณการจัดส่ง** เริ่มจากการรวบรวมพิกัด (ตำแหน่ง) ของลูกค้าและศูนย์กระจายสินค้าปัจจุบัน และปริมาณการจัดส่งสินค้าหรือความต้องการของแต่ละจุด
2. **คำนวณตำแหน่งศูนย์กลางความโน้มถ่วง** สามารถใช้สูตรทางคณิตศาสตร์เพื่อคำนวณตำแหน่งศูนย์กลางความโน้มถ่วง โดยคำนวณจากค่าเฉลี่ยของพิกัดที่ถ่วงน้ำหนักด้วยปริมาณการจัดส่ง

$$X_{cg} = \frac{\sum_{i=1}^n x_i w_i}{\sum_{i=1}^n w_i} \quad (7.1)$$

$$Y_{cg} = \frac{\sum_{i=1}^n y_i w_i}{\sum_{i=1}^n w_i} \quad (7.2)$$

โดยที่:

X_{cg} และ Y_{cg} คือพิกัดของศูนย์กลางความโน้มถ่วง

x_i, y_i คือพิกัดของจุดที่ i

w_i คือปริมาณการจัดส่งสินค้าหรือความต้องการของจุดที่ i

n คือจำนวนจุดทั้งหมด

3. การประเมินและปรับแต่ง

หลังจากคำนวณตำแหน่งศูนย์กลางความโน้มถ่วงแล้ว ให้ประเมินว่าตำแหน่งนั้นเป็นตำแหน่งที่เหมาะสมที่สุดหรือไม่ โดยพิจารณาปัจจัยเพิ่มเติมเช่น โครงสร้างพื้นฐาน การเข้าถึงการขนส่ง และกฎระเบียบท้องถิ่น ปรับแต่งตำแหน่งให้เหมาะสมยิ่งขึ้นหากจำเป็น ข้อดีของวิธีศูนย์กลางความโน้มถ่วงประกอบไปด้วย 2 ได้แก่ วิธีนี้ใช้สูตรทางคณิตศาสตร์ที่ไม่นับซับซ้อนและสามารถคำนวณได้ง่าย และวิธีนี้พิจารณาปริมาณการจัดส่งสินค้าหรือความต้องการของแต่ละจุด ทำให้ตำแหน่งที่คำนวณได้มีความแม่นยำและเหมาะสม อย่างไรก็ตามวิธีนี้ก็ยังมีข้อจำกัดได้แก่ วิธีนี้ไม่ได้พิจารณาถึงข้อจำกัดทางกายภาพเช่น การเข้าถึงถนน สภาพแวดล้อมท้องถิ่น หรือกฎระเบียบทางกฎหมาย และ วิธีนี้สมมติว่าการกระจายตัวของปริมาณการจัดส่งหรือความต้องการมีค่าสม่ำเสมอ ซึ่งอาจไม่เป็นจริงในสถานการณ์จริง

วิธีศูนย์กลางความโน้มถ่วงเป็นเครื่องมือที่มีประโยชน์สำหรับการตัดสินใจเกี่ยวกับตำแหน่งที่ตั้งในซัพพลายเชน แต่ควรใช้งานร่วมกับการวิเคราะห์เพิ่มเติมเพื่อให้ได้ผลลัพธ์ที่ดีที่สุด (Lukardi & Hamsal, 2020; Şahin, 2021)

ตัวอย่าง 7.1 การหาที่ตั้งด้วยวิธีศูนย์กลางความโน้มถ่วง

บริษัท XYZ ต้องการหาตำแหน่งที่ตั้งที่เหมาะสมสำหรับคลังสินค้าแห่งใหม่ โดยใช้วิธี Center of Gravity Method โดยมีลูกค้า 4 รายในพื้นที่ที่ต้องการให้บริการดังนี้:

ตารางที่ 7.3 ข้อมูลพิกัด (X, Y)

ลูกค้า	พิกัด X (กิโลเมตร)	พิกัด Y (กิโลเมตร)	ปริมาณการขนส่ง (ตันต่อปี)
A	10	20	50
B	20	30	30
C	25	25	20
D	30	10	40

ที่มา: จัดทำโดยผู้แต่ง

วิธีคำนวณ:

1. คำนวณตำแหน่งที่ตั้งของคลังสินค้า (X, Y) โดยใช้สูตร Center of Gravity Method: ตามสมการที่ 7.1 และ 7.2
2. หาค่า X_{cg} แทนค่าลงในสมการที่ 7.1 ได้ดังต่อไปนี้

$$X_{cg} = \frac{(10 \times 50) + (20 \times 30) + (25 \times 20) + (30 \times 40)}{50 + 30 + 20 + 40}$$

$$X_{cg} = \frac{(500) + (600) + (500) + (1200)}{140}$$

$$X_{cg} = \frac{2800}{140} = 20 \text{ km.}$$

3. หาค่า Y_{cg} แทนค่าลงในสมการที่ 7.2 ได้ดังต่อไปนี้

$$Y_{cg} = \frac{(20 \times 50) + (30 \times 30) + (25 \times 20) + (10 \times 40)}{50 + 30 + 20 + 40}$$

$$Y_{cg} = \frac{(1000) + (900) + (500) + (400)}{140}$$

$$Y_{cg} = \frac{2800}{140} = 20 \text{ km.}$$

4. ตำแหน่งที่ตั้งคลังสินค้าที่เหมาะสมที่สุดตามวิธี Center of Gravity คือที่พิกัด $X_{cg} = 20$ กิโลเมตร และ $Y_{cg} = 20$ กิโลเมตร

7.4.3 การรวมศูนย์ (Centralization) vs การกระจายศูนย์ (De-centralization)

การตัดสินใจในการรวมศูนย์หรือกระจายศูนย์ในซัพพลายเชนและการจัดการโลจิสติกส์มีข้อดีและข้อเสียที่ต้องพิจารณา การเลือกแนวทางที่เหมาะสมขึ้นอยู่กับปัจจัยต่าง ๆ เช่น ประเภทของสินค้า ความต้องการของลูกค้า ต้นทุน และโครงสร้างพื้นฐาน

การรวมศูนย์ (Centralization) มีข้อดีหลายประการได้แก่ ประหยัดต้นทุน (cost savings) เพราะลดต้นทุนการดำเนินงานโดยการรวมคลังสินค้าและศูนย์กระจายสินค้าให้น้อยลง และลดต้นทุนการจัดการสินค้าคงคลังเนื่องจากสามารถควบคุมได้จากจุดศูนย์กลางเดียว (Cachon & Fisher, 2000) สามารถควบคุมและบริหารจัดการที่ดีขึ้น ทำให้สามารถควบคุมการดำเนินงานทั้งหมดได้ง่ายขึ้น และการบริหารจัดการสต็อกและการวางแผนการจัดส่งสามารถทำได้อย่างมีประสิทธิภาพมากขึ้น ช่วยให้ใช้ทรัพยากรอย่างมีประสิทธิภาพ สามารถใช้ทรัพยากรที่มีอยู่ให้เกิดประโยชน์สูงสุด เช่น อุปกรณ์และบุคลากร

อย่างไรก็ตามการรวมศูนย์ก็มีข้อเสียได้แก่ ความเสี่ยงในการหยุดชะงักหากเกิดปัญหาที่ศูนย์กลาง เช่น ภัยพิบัติ หรือปัญหาทางเทคนิค จะส่งผลกระทบต่อซัพพลายเชนทั้งหมด ความยืดหยุ่นต่ำทำให้การตอบสนองต่อความต้องการของตลาดที่เปลี่ยนแปลงอาจทำได้ช้ากว่า และการปรับเปลี่ยนการดำเนินงานเพื่อรองรับความต้องการเฉพาะของลูกค้าอาจทำได้ยาก ระยะเวลาการจัดส่งยาวขึ้นทำให้ลูกค้าที่อยู่ไกลจากศูนย์กลางอาจต้องใช้เวลารอคอยสินค้านานขึ้น

การกระจายศูนย์ (De-centralization) มีข้อดีหลายประการได้แก่ ความยืดหยุ่นสูงทำให้สามารถตอบสนองต่อความต้องการของตลาดที่เปลี่ยนแปลงได้อย่างรวดเร็ว และสามารถปรับเปลี่ยนการดำเนินงานเพื่อรองรับความต้องการเฉพาะของลูกค้าได้ง่ายกว่า ลดความเสี่ยงในการหยุดชะงักหากเกิดปัญหาที่ศูนย์ใดศูนย์หนึ่ง จะไม่ส่งผลกระทบต่อซัพพลายเชนทั้งหมด (Ivanov & Dolgui, 2020) ระยะเวลาการจัดส่งสั้นลงเพราะการกระจายศูนย์กระจายสินค้าไปอยู่ใกล้ลูกค้าทำให้ได้รับสินค้ารวดเร็วยิ่งขึ้น เช่นเดียวกัน การกระจายศูนย์ย่อมมีข้อเสีย ได้แก่ ต้นทุนการดำเนินงานสูงขึ้น เพราะการบริหารจัดการหลายศูนย์กระจายสินค้าอาจทำให้ต้นทุนการดำเนินงานสูงขึ้น รวมไปถึงการจัดการสินค้าคงคลังในหลายจุดอาจทำได้ยากขึ้นและมีค่าใช้จ่ายสูง การควบคุมและบริหารจัดการที่ซับซ้อน เพราะการบริหารจัดการศูนย์กระจายสินค้าหลายแห่งอาจทำให้การควบคุมเป็นไปได้ยากขึ้น นอกจากนี้การวางแผนและการประสานงานต้องใช้ความพยายามมากขึ้น รวมไปถึงการใช้ทรัพยากรที่ไม่มีประสิทธิภาพเพราะทรัพยากรอาจถูกใช้งานไม่เต็มประสิทธิภาพเนื่องจากต้องกระจายไปยังหลายศูนย์

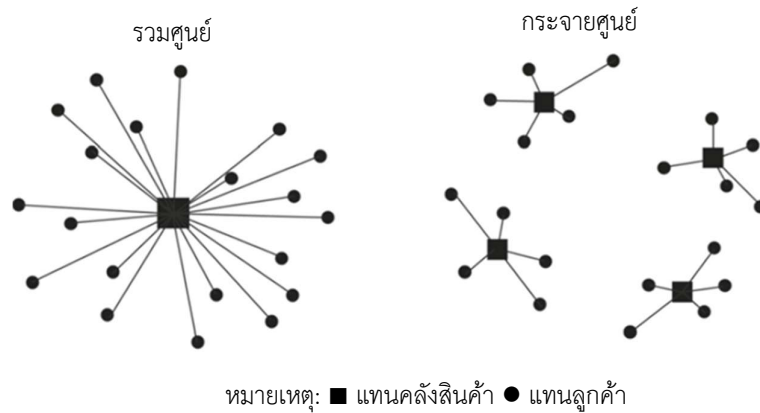
7.4.4 กฎรากที่สอง (Square Root Law, SRL)

กฎรากที่สอง (square root law) เป็นแนวคิดที่ใช้ในการจัดการสินค้าคงคลังและการกระจายศูนย์คลังสินค้าในซัพพลายเชน กฎนี้ช่วยให้บริษัทสามารถประมาณจำนวนคลังสินค้าที่

เหมาะสมเพื่อให้บรรลุวัตถุประสงค์ด้านต้นทุนและระดับการบริการ โดยอิงจากการเปลี่ยนแปลงของจำนวนศูนย์คลังสินค้า

สูตรกฎรากที่สอง

กฎรากที่สองแสดงความสัมพันธ์ระหว่างจำนวนคลังสินค้าที่มีอยู่กับปริมาณสินค้าคงคลังที่จำเป็น:



ภาพที่ 7.1 การรวมศูนย์และการกระจายศูนย์
ที่มา: ดัดแปลงจาก (Liu, 2022)

$$S_2 = S_1 \times \sqrt{n_2/n_1} \quad (7.3)$$

โดยที่

S_2 คือ ปริมาณสินค้าคงคลังทั้งหมดที่ต้องการเมื่อมีคลังสินค้าที่ใหม่

S_1 คือ ปริมาณสินค้าคงคลังทั้งหมดที่ต้องการเมื่อมีคลังสินค้าที่เดิม

n_1 คือ จำนวนคลังสินค้าที่เดิม

n_2 คือ จำนวนคลังสินค้าที่ใหม่

การประยุกต์ใช้กฎรากที่สอง

กฎรากที่ 2 จะประมาณต้นทุนสินค้าคงคลังเมื่อเพิ่มจำนวนคลังสินค้า จำนวนสินค้าคงคลังที่ต้องการจะเพิ่มขึ้นตามรากที่สองของจำนวนคลังสินค้าที่เพิ่มขึ้น ไม่ใช่เพิ่มขึ้นตามสัดส่วนตรง เช่นเดียวกับการเพิ่มจำนวนคลังสินค้าจะเพิ่มระดับการบริการเพราะการเพิ่มจำนวนคลังสินค้า

สามารถช่วยลดระยะเวลาการจัดส่งและปรับปรุงการตอบสนองต่อความต้องการของลูกค้า แต่ก็จะต้องคำนึงถึงการเพิ่มขึ้นของต้นทุนสินค้าคงคลังประกอบด้วย นอกจากนี้กฎรากที่สองช่วยในการวางแผนเครือข่ายคลังสินค้า ทำให้สามารถคำนวณและปรับปรุงจำนวนคลังสินค้าเพื่อให้เหมาะสมกับเป้าหมายทางธุรกิจได้ ข้อดีของการใช้กฎรากที่สองมีหลายประการ หลัก ๆ ก็คือ ความง่ายในการคำนวณเพราะสูตรที่ไม่ซับซ้อน ช่วยทำให้ปรับสมดุลระหว่างต้นทุนสินค้าคงคลังและระดับการบริการที่เหมาะสม อย่างไรก็ตามการใช้กฎรากที่สองก็มีข้อพึงระวังได้แก่ สมมติฐานของสูตรคณิตศาสตร์ทำให้ไม่เหมาะสมที่จะใช้กับทุกสถานการณ์ รวมไปถึง สูตรนี้ไม่คำนึงถึงปัจจัยเช่น ความผันแปรของความต้องการ การเปลี่ยนแปลงของตลาด หรือปัจจัยอื่น ๆ ที่มีผลต่อสินค้าคงคลัง

ตัวอย่าง 7.2 กฎรากที่สอง

สมมติว่าเดิมบริษัทมีคลังสินค้า 4 แห่ง และต้องการสินค้าคงคลังรวม 2000 หน่วย และบริษัทกำลังพิจารณาเพิ่มจำนวนคลังสินค้าเป็น 9 แห่ง สามารถคำนวณปริมาณสินค้าคงคลังที่ต้องการใหม่ได้ดังนี้:

$$S_2 = 2000 \times \sqrt{9/4} = 3000 \text{ หน่วย}$$

ดังนั้น หากบริษัทเพิ่มจำนวนคลังสินค้าเป็น 9 แห่ง ปริมาณสินค้าคงคลังรวมที่ต้องการจะเพิ่มขึ้นเป็น 3000 หน่วย โดยสรุปภายใต้สมมติฐานของกฎรากที่สอง (square root law - SRL) การรวมศูนย์ (centralization) จะลดระดับสินค้าคงคลังโดยอัตโนมัติ แม้ว่าในทางปฏิบัติการใช้งานของกฎนี้อาจยังมีข้อถกเถียงอยู่ (Oeser & Romano, 2016) แต่ SRL อธิบายความสัมพันธ์ระหว่างสินค้าคงคลังในคลังสินค้านรวมศูนย์และกระจายศูนย์ ซึ่งให้ข้อมูลเชิงลึกที่มีประโยชน์ในการตัดสินใจเชิงบริหารในการออกแบบเครือข่ายโลจิสติกส์

7.5 การออกแบบเครือข่ายโลจิสติกส์โดยใช้แบบจำลองเชิงเส้น (Linear Programming, LP)

ในหัวข้อนี้จะอธิบายโดยใช้ตัวอย่างที่ 7.3 - 7.5 เป็นกรณีศึกษาของธุรกิจแฟรนไชส์พิซซ่าที่ตัดสินใจเปิดร้านเพิ่มเติมเพื่อครอบคลุมความต้องการในพื้นที่ที่แตกต่างกันภายใต้ข้อจำกัดบางประการโดยมีสรุปข้อมูลเบื้องต้นดังนี้

ตัวอย่าง 7.3 การออกแบบเครือข่ายโลจิสติกส์โดยใช้แบบจำลองเชิงเส้น

ข้อมูลเบื้องต้น พบว่า มีห้าพื้นที่ที่ต้องการเปิดร้านใหม่ แต่ละพื้นที่ที่มีความต้องการพิซซ่าที่แตกต่างกัน มีข้อจำกัดในการเปิดร้าน เช่น ต้นทุน งบประมาณ และจำนวนร้านสูงสุดที่สามารถเปิดได้ โดยมีจุดมุ่งหมายคือการครอบคลุมความต้องการสูงสุดภายใต้ข้อจำกัดเหล่านี้

ขั้นตอนในการแก้ปัญหา แสดงดังต่อไปนี้

1. **การรวบรวมข้อมูล** เริ่มจากการระบุตำแหน่งของห้าพื้นที่ดังกล่าว ระบุความต้องการพิซซ่าในแต่ละพื้นที่รวมถึงระบุต้นทุนในการเปิดร้านในแต่ละพื้นที่ หลังจากนั้นทำการระบุข้อจำกัด เช่น งบประมาณสูงสุด จำนวนร้านสูงสุดที่สามารถเปิดได้
2. **การกำหนดรูปแบบปัญหา** เริ่มจากกำหนดเทคนิคการเพิ่มประสิทธิภาพที่จะใช้ เช่น การเขียนโปรแกรมเชิงเส้น (linear programming) หรือการเขียนโปรแกรมเชิงจำนวนเต็ม (integer programming) เพื่อแก้ปัญหา โดยจะต้องมีการกำหนดตัวแปรที่เกี่ยวข้อง เช่น x_i คือ จำนวนร้านที่จะเปิดในพื้นที่ i หลังจากมีตัวแปรแล้วจึงกำหนดฟังก์ชันวัตถุประสงค์ (objective function) เช่น การเพิ่มการครอบคลุมความต้องการ สุดท้ายจึงทำการกำหนดข้อจำกัด (constraints) เช่น ต้นทุนรวมต้องไม่เกินงบประมาณ จำนวนร้านที่เปิดต้องไม่เกินข้อจำกัด
3. **การแก้ปัญหา** จะใช้ซอฟต์แวร์หรือเครื่องมือทางคณิตศาสตร์ เช่น Excel Solver, Python, หรือซอฟต์แวร์การแก้ปัญหาทางคณิตศาสตร์อื่น ๆ เพื่อหาคำตอบที่เหมาะสมที่สุด

สมมติว่า:

- 1) มีพื้นที่ 5 แห่งที่ต้องการเปิดร้านพิซซ่า: A, B, C, D, และ E
- 2) ความต้องการในแต่ละพื้นที่คือ: 100, 150, 200, 250, และ 300 หน่วยตามลำดับ

- 3) ต้นทุนในการเปิดร้านในแต่ละพื้นที่คือ: 50,000, 60,000, 55,000, 70,000, และ 65,000 บาทตามลำดับ
- 4) งบประมาณรวมคือ 200,000 บาท
- 5) จำนวนร้านสูงสุดที่สามารถเปิดได้คือ 4 ร้าน

ฟังก์ชันวัตถุประสงค์และข้อจำกัด

Objective function: Maximize $100x_A + 150x_B + 200x_C + 250x_D + 300x_E$

Subject to:

$$50,000x_A + 60,000x_B + 55,000x_C + 70,000x_D + 65,000x_E \leq 200,000$$

$$x_A + x_B + x_C + x_D + x_E \leq 4$$

$$x_i \in \{0,1\} \quad \forall i \in \{A, B, C, D, E\} \text{ (หมายความว่าแต่ละพื้นที่สามารถเปิดร้านได้หรือไม่ได้)}$$

ตัวอย่าง 7.4 ปัญหาการออกแบบเครือข่ายโลจิสติกส์สำหรับร้านพิซซ่าแบบมีข้อจำกัดความจุ

ปัญหาการออกแบบโครงข่ายโลจิสติกส์ข้าม 5 ภูมิภาค ตัวอย่างเช่น R1 R2 R3 R4 R5 ร้านพิซซ่ากำลังต้องการจะเปิดร้านเพื่อตอบสนองความต้องการใน 5 ภูมิภาคนี้ ดังนั้นในฐานะผู้จัดการโลจิสติกส์ คุณได้รับความต้องการสินค้ารายวันของแต่ละภูมิภาคนี้ ต้นทุนค่าขนส่งเฉลี่ยระหว่าง 2 ภูมิภาค และต้นทุนคงที่ในการเปิดร้านในแต่ละภูมิภาค โดยที่ร้านที่จะเปิดเพิ่มนี้จะมีการออกแบบที่เหมือนกันอยู่ 2 ทางเลือก คือ ความจุสูง (High Capacity) และความจุต่ำ (Low Capacity) โดยมีวัตถุประสงค์เพื่อทำให้ต้นทุนรวมต่ำที่สุดและเพียงพอที่จะสามารถตอบสนองความต้องการของทั้ง 5 ภูมิภาคนี้

ปัญหานี้เป็นปัญหาการทำต้นทุนให้ต่ำที่สุด โดยสามารถสร้างแบบจำลองเชิงเส้นได้ดังต่อไปนี้

$$\min \sum_{i=1}^n f_i y_i + \sum_{i=1}^n \sum_{j=1}^m d_{ij} x_{ij} \quad (7.4)$$

โดยที่

f_i = ต้นทุนคงที่ถ้าจะเปิดร้าน i

y_i = ตัวแปรฐานสอง $y_i = 1$ ถ้าร้านเปิด แต่เป็น 0 ร้านปิด

d_{ij} = ต้นทุนค่าขนส่งเฉลี่ยต่อหน่วยจากร้าน i ไปยังภูมิภาค j

x_{ij} = ปริมาณที่ขนส่งจากร้าน i ไปยังภูมิภาค j

n = จำนวนร้านทั้งหมดที่จะเปิด

m = จำนวนภูมิภาค

โดยมีข้อจำกัดดังต่อไปนี้

$$\sum_{i=1}^n x_{ij} = D_j \quad \text{สำหรับ } j = 1, \dots, m \quad (7.5)$$

$$\sum_{j=1}^m x_{ij} = C_i \quad \text{สำหรับ } i = 1, \dots, m \quad (7.6)$$

โดยที่

D_j เป็นความต้องการทั้งหมดของภูมิภาค j

C_i เป็นความจุของร้าน i

ข้อจำกัดในสมการ (7.5) กำหนดให้ปริมาณทั้งหมดที่จัดส่งจากร้านที่เป็นไปได้ต้องตอบสนองความต้องการทั้งหมดจากแต่ละภูมิภาค และ ข้อจำกัดในสมการ (7.6) ข้อจำกัดนี้กำหนดให้ปริมาณทั้งหมดที่จัดส่งจากร้านไม่เกินความจุที่เป็นไปได้ของร้านนั้น (ทั้งความจุสูงหรือความจุต่ำ):

ตารางที่ 7.4 ตารางข้อมูลสำหรับตัวอย่างที่ 7.4

ร้าน	ความต้องการ สินค้า	ต้นทุนคงที่		ทางเลือกความจุ		ต้นทุนค่าขนส่ง				
		ต่ำ	สูง	ต่ำ	สูง	R1	R2	R3	R4	R5
1(R1)	5000	800	1200	2000	6000	5	12	17	8	7
2(R2)	1800	1250	1650	2000	6000	12	4	10	7	8
3(R3)	800	900	1450	2000	6000	17	10	3	11	8
4(R4)	630	650	800	2000	6000	8	7	11	7	18
5(R5)	490	666	950	2000	6000	7	8	10	18	6

ที่มา: จัดทำโดยผู้แต่ง

	Delivery	Amount		loc	Low_cap	High_cap
0	R1 to R1	5000.0				
1	R1 to R2	0.0	0	R1	0.0	1.0
2	R1 to R3	0.0	1	R2	0.0	1.0
3	R1 to R4	0.0	2	R3	1.0	0.0
4	R1 to R5	490.0	3	R4	0.0	0.0
5	R2 to R1	0.0	4	R5	0.0	0.0
6	R2 to R2	1800.0				
7	R2 to R3	0.0				
8	R2 to R4	630.0				
9	R2 to R5	0.0				
10	R3 to R1	0.0	Total Optimized Costs = £46190.0			
11	R3 to R2	0.0				
12	R3 to R3	800.0				
13	R3 to R4	0.0				
14	R3 to R5	0.0				
15	R4 to R1	0.0				
16	R4 to R2	0.0				
17	R4 to R3	0.0				
18	R4 to R4	0.0				
19	R4 to R5	0.0				
20	R5 to R1	0.0				
21	R5 to R2	0.0				
22	R5 to R3	0.0				
23	R5 to R4	0.0				
24	R5 to R5	0.0				

ภาพที่ 7.2 ผลลัพธ์จากการแก้ปัญหา LP โดยใช้โปรแกรม Python

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลจาก โปรแกรม Python)

จากผลลัพธ์เราจะเห็นได้ว่าร้านพิซซ่าใหม่ควรเปิดในเขต R1, R2 และ R3 เท่านั้น และร้านใน R1 และ R2 ควรมีการออกแบบความจุสูงแม้ว่าจะมีความจุเหลือเฟือก็ตาม ในขณะที่ร้านใน R3 สามารถมีตัวเลือกความจุต่ำได้

โดยเฉพาะอย่างยิ่ง ร้านใน R1 ควรครอบคลุมทั้งเขต R1 และ R5 และร้านใน R2 ควรครอบคลุมทั้งเขต R2 และ R4 ส่วนร้านใน R3 จะให้บริการเฉพาะเขตของตัวเองเท่านั้น ตามทางออกที่ได้รับการปรับให้เหมาะสมแล้ว ต้นทุนรวมจะอยู่ที่ £46,190 หรือ 1,719,353 บาท

ตัวอย่าง 7.5 ปัญหาการเลือกตำแหน่งที่ตั้งอย่างง่ายและแก้ปัญหาด้วย Microsoft Excel Solver

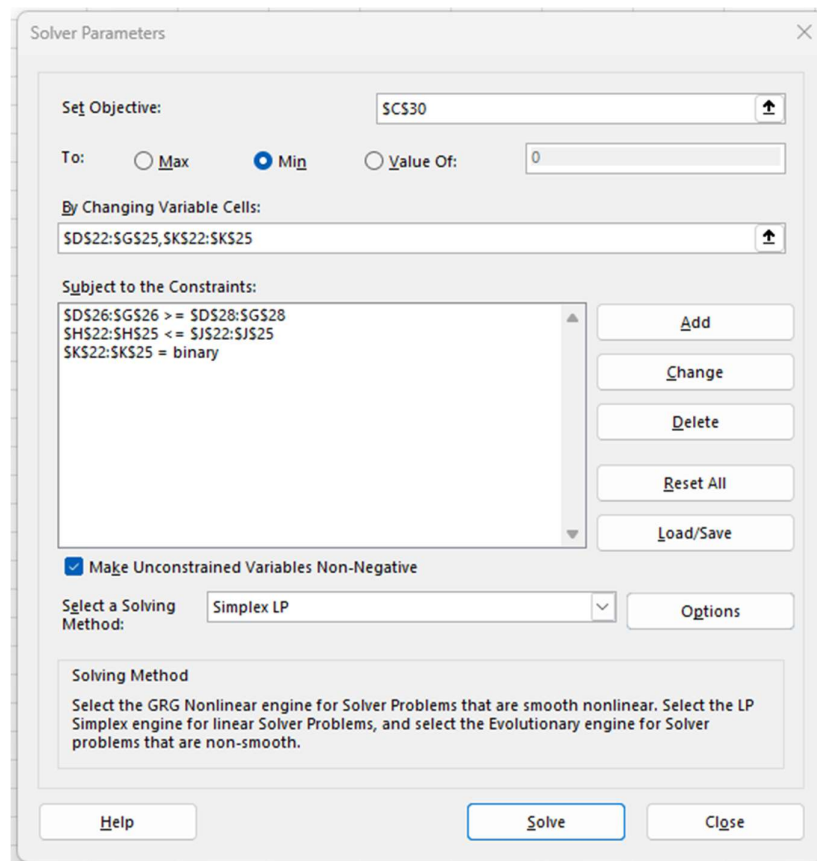
ในตัวอย่างนี้พยายามหาต้นทุนต่ำสุดในการขนส่งคอมพิวเตอร์จากโรงงานผลิต (ตั้งอยู่ที่ซานฟรานซิสโก ลอสแอนเจลิส ฟีนิกซ์ หรือเดนเวอร์) ไปยังร้านค้า (ซานดีเอโก บาร์สตอว์ ทูซอน ดัลลัส) โดยกำหนดว่าจะสร้างโรงงานที่ใดเพื่อตอบสนองความต้องการของร้านค้าต่างๆ

	A	B	C	D	E	F	G	H	I	J	K	L
1												
2			Facility Location Example									
3			<u>Description:</u>									
4			The facility location problem seeks to minimise or maximise the cost of a set of flows along arcs directly from origins									
5			to destinations by determining where to set up facilities while being restricted by available supplies, required									
6			demands, and bounds on arc flows.									
7			<u>Example context:</u>									
8			This model seeks a minimum cost of transporting computers from production plants (located at San Francisco, Los									
9			Angeles, Phoenix, or Denver) to stores (San Diego, Barstow, Tucson, Dallas) by determining where to build the plants									
10			while satisfying store demands.									
11												
12			<u>Model:</u>									
13			<u>Costs</u>									
14				San Diego	Barstow	Tucson	Dallas	Fixed				
15			San Francisco	5	3	2	6	70				
16			Los Angeles	10	7	8	9	70				
17			Phoenix	6	5	3	8	65				
18			Denver	9	8	6	5	70				
19												
20			<u>Constraints and solution</u>									
21				San Diego	Barstow	Tucson	Dallas	VALUE		Supply	Build?	
22			San Francisco	700	1000	0	0	1700	<=	1700	^b	1
23			Los Angeles	0	0	0	0	0	<=	0	^b	0
24			Phoenix	200	0	1500	0	1700	<=	1700	^b	1
25			Denver	800	0	0	1200	2000	<=	2000	^b	1
26			VALUE	1700	1000	1500	1200					
27				>=	>=	>=	>=					
28			Demand	≥	1700	1000	1500	1200				
29			min									
30			Total cost	25605								

ภาพที่ 7.3 โมเดลเชิงเส้นของตัวอย่างที่ 7.5

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลจาก โปรแกรม MS Excel)

เราสามารถแก้ปัญหาของโมเดลได้โดยป้อนข้อมูลลงใน Excel Solver ดังต่อไปนี้



ภาพที่ 7.4 การป้อนข้อมูลลงใน Excel Solver ตัวอย่างที่ 7.5

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลจาก โปรแกรม MS Excel)

หลังจากนั้นเมื่อกด Solve จะได้ผลลัพธ์คือ

ส่งสินค้าจาก SF → SD จำนวน 700 หน่วย

ส่งสินค้าจาก SF → BT จำนวน 1000 หน่วย

ส่งสินค้าจาก PN → SD จำนวน 200 หน่วย

ส่งสินค้าจาก PN → TS จำนวน 1500 หน่วย

ส่งสินค้าจาก DV → SD จำนวน 800 หน่วย

ส่งสินค้าจาก DV → DL จำนวน 1200 หน่วย

เลือกสร้างโรงงานที่ SF, PN, DV

ต้นทุนรวม = \$25605 หรือ 857,716 บาท

สรุป

การจัดการโลจิสติกส์ (Logistics Management) คือกระบวนการวางแผน ดำเนินการ และควบคุมการเคลื่อนย้ายและจัดเก็บสินค้า วัตถุดิบ ข้อมูล และทรัพยากรต่าง ๆ จากจุดเริ่มต้นถึงปลายทาง โดยมุ่งเน้นให้เกิดความคุ้มค่าในเชิงต้นทุนและเวลา รวมถึงการตอบสนองต่อความต้องการของลูกค้าได้อย่างมีประสิทธิภาพ การจัดการโลจิสติกส์ครอบคลุมการจัดการสินค้าคงคลัง การขนส่ง การจัดเก็บสินค้า การดำเนินการคลังสินค้า และการจัดการข้อมูลและระบบสารสนเทศที่เกี่ยวข้อง การจัดการโลจิสติกส์มีความสำคัญในหลายๆด้าน ได้แก่ (1) เพิ่มประสิทธิภาพการดำเนินงาน (2) ช่วยลดต้นทุนด้านต่างๆ (3) ปรับปรุงการบริการลูกค้า (4) ช่วยจัดการความเสี่ยง (5) สนับสนุนการตัดสินใจที่มี (6) สนับสนุนการขยายธุรกิจ การจัดการโลจิสติกส์ที่ดีประกอบด้วย การเลือกโหมดการขนส่งที่เหมาะสม และการออกแบบโครงข่ายโลจิสติกส์ที่มีประสิทธิภาพ โดยโหมดการขนส่งประกอบไปด้วย 5 โหมดได้แก่ การขนส่ง ทางบก ทางน้ำ ทางอากาศ ทางราง และทางท่อ การเลือกโหมดการขนส่งที่เหมาะสมขึ้นอยู่กับหลายปัจจัย เช่น ลักษณะของสินค้า ระยะทาง ต้นทุน เวลาที่ต้องใช้ในการขนส่ง และความต้องการของลูกค้า โดยการวิเคราะห์และเลือกโหมดการขนส่งที่เหมาะสมจะช่วยเพิ่มประสิทธิภาพในซัพพลายเชนและลดต้นทุนโลจิสติกส์โดยรวม สำหรับการออกแบบโครงข่ายโลจิสติกส์ที่มีประสิทธิภาพเริ่มจากการหาทำเลที่ตั้งที่เหมาะสมจากการคำนึงถึงหลายปัจจัยได้แก่ ต้นทุนที่เกี่ยวข้อง แรงงานที่มีทักษะ กฎหมายและสิ่งแวดล้อม ความใกล้เคียงแหล่งวัตถุดิบและตลาด และปัจจัยอื่นๆ รวมไปถึงการวางกลยุทธ์โครงข่ายโลจิสติกส์ว่าต้องการจะไปในทางเลือกระหว่าง การตอบสนองต่อลูกค้าที่มีความยืดหยุ่นและรวดเร็ว หรือการมุ่งเน้นความเป็นเลิศของผลิตภาพ (Productivity) ในการใช้ทรัพยากรที่น้อยลง ซึ่งการวางกลยุทธ์นี้เองจะกำหนดให้โครงข่ายโลจิสติกส์ให้เป็นแบบ กระจายศูนย์ (เน้นการตอบสนองต่อความต้องการของลูกค้าที่รวดเร็วและยืดหยุ่น) หรือรวมศูนย์ (เน้นความเป็นเลิศทางด้านผลิตภาพ) สอดคล้องตามกัน นอกจากนี้การออกแบบโครงข่ายโลจิสติกส์ยังสามารถใช้เครื่องมือทางคณิตศาสตร์เข้ามาช่วยในการออกแบบอย่างเช่น การใช้แบบจำลองเชิงเส้นตรง (Linear Programming) โดยผลลัพธ์ที่ได้จะช่วยให้สามารถกำหนดเส้นทางการขนส่ง หรือ ทำเลที่เหมาะสมในการเปิดร้านค้าหรือศูนย์กระจายสินค้า โดยที่มีต้นทุนรวมที่เกี่ยวข้องต่ำที่สุด

แบบฝึกหัด

1. อธิบายความสำคัญของการจัดการโลจิสติกส์ในซัพพลายเชน และผลกระทบที่อาจเกิดขึ้นหากการจัดการโลจิสติกส์ไม่มีประสิทธิภาพ
2. จากข้อมูลด้านล่างนี้ ให้นักเรียนคำนวณต้นทุนรวมในการขนส่งสินค้าจากจุด A ไปยังจุด B โดยเปรียบเทียบระหว่างการขนส่งทางถนนและทางราง และวิเคราะห์ว่าโหมดการขนส่งใดที่เหมาะสมที่สุด

ข้อมูล:

น้ำหนักสินค้า: 20 ตัน

ระยะทาง: 500 กิโลเมตร

ต้นทุนการขนส่งทางถนน: 2 บาท/กิโลเมตร/ตัน

ต้นทุนการขนส่งทางราง: 1.5 บาท/กิโลเมตร/ตัน

เวลาในการขนส่ง: ทางถนน 8 ชั่วโมง, ทางราง 12 ชั่วโมง

3. บริษัท ABC มีสินค้าคงคลังที่ต้องจัดการอย่างมีประสิทธิภาพ แต่พบว่าสินค้าคงคลังจำนวนมากในคลังสินค้า ให้นักเรียนวิเคราะห์ปัญหาที่อาจเกิดขึ้นจากการบริหารสินค้าคงคลังที่ไม่เหมาะสม และเสนอแนวทางในการปรับปรุงการจัดการสินค้าคงคลัง
4. ในกรณีที่บริษัทต้องการขยายธุรกิจไปยังตลาดต่างประเทศ ให้นักเรียนเลือกโหมดการขนส่งที่เหมาะสมที่สุดสำหรับการขนส่งสินค้าไปยังประเทศที่อยู่ห่างไกล พร้อมอธิบายเหตุผลในการเลือกโหมดการขนส่งดังกล่าว
5. บริษัท ABC มีศูนย์กระจายสินค้าจำนวน 4 แห่ง โดยมีสินค้าคงคลังรวมทั้งหมด 10,000 หน่วยในปัจจุบัน หากบริษัทต้องการลดจำนวนศูนย์กระจายสินค้าเหลือ 2 แห่ง คาดว่าจะต้องมีสินค้าคงคลังรวมเท่าไรตาม Square Root Law
6. บริษัท ABC ต้องการหาทำเลที่ตั้งคลังสินค้าแห่งใหม่ที่จะสามารถให้บริการกับลูกค้า 3 แห่งในพื้นที่ใกล้เคียง โดยมีข้อมูลตำแหน่ง (X, Y) และปริมาณการส่งสินค้า (น้ำหนัก) ของแต่ละลูกค้าดังนี้:

ลูกค้า	ตำแหน่ง X (กิโลเมตร)	ตำแหน่ง Y (กิโลเมตร)	ปริมาณการส่ง สินค้า (ตัน)
ลูกค้า A	10	20	50
ลูกค้า B	30	40	30
ลูกค้า C	50	60	20

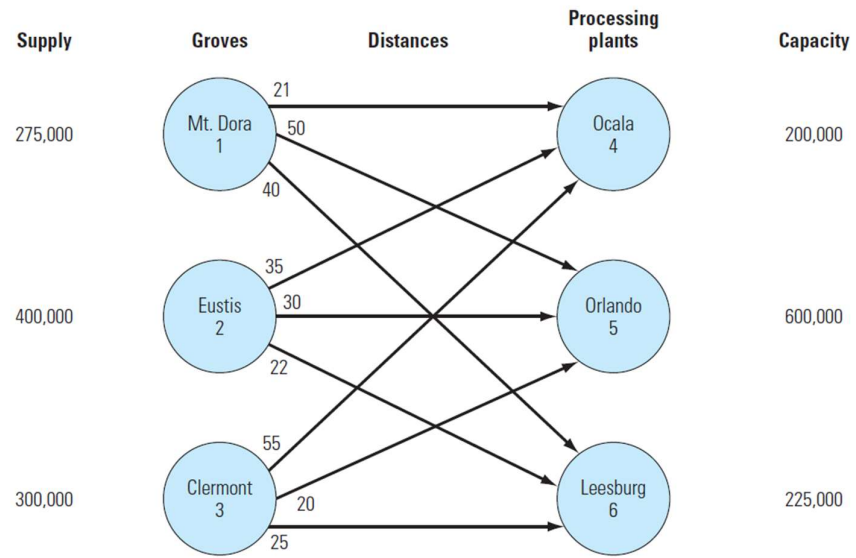
ให้นักศึกษาคำนวณหาตำแหน่งที่ตั้งคลังสินค้าที่เหมาะสมโดยใช้วิธี Center of Gravity Method และแสดงวิธีคำนวณอย่างละเอียด

- บริษัทต้องการใช้ Center of Gravity Method เพื่อหาทำเลที่ตั้งโรงงานใหม่ แต่มูลค่าการขนส่งของแต่ละลูกค้าไม่เท่ากันและตำแหน่งของลูกค้ามีการเปลี่ยนแปลงบ่อยครั้ง ให้นักเรียนอธิบายถึงการปรับปรุงวิธี Center of Gravity Method เพื่อให้เหมาะสมกับสถานการณ์นี้
- Tropicsun** เป็นผู้นำในการปลูกและจัดจำหน่ายผลิตภัณฑ์ผลไม้รสไม่รสจัด โดยมีสวนส้มขนาดใหญ่สามแห่งที่กระจายอยู่ทั่วตอนกลางของฟลอริดาในเมือง Mt. Dora, Eustis, และ Clermont ปัจจุบัน Tropicsun มีส้ม 275,000 บุชเชลที่สวนใน Mt. Dora, 400,000 บุชเชลที่สวนใน Eustis, และ 300,000 บุชเชลที่สวนใน Clermont

Tropicsun มีโรงงานแปรรูปส้มใน Ocala, Orlando, และ Leesburg โดยมีความสามารถในการแปรรูปส้มได้ 200,000, 600,000 และ 225,000 บุชเชลตามลำดับ

Tropicsun ทำสัญญากับบริษัทขนส่งในท้องถิ่นเพื่อขนส่งผลไม้จากสวนไปยังโรงงานแปรรูป บริษัทขนส่งจะคิดอัตราค่าขนส่งคงที่ต่อไมล์สำหรับทุกบุชเชลของผลไม้ที่ต้องขนส่ง แต่ละไมล์ที่บุชเชลของผลไม้เดินทางไปเรียกว่า bushel-mile สรุประยะทาง (เป็นไมล์) ระหว่างสวนและโรงงานแปรรูปแสดงดังรูปข้างล่าง:

ปัญหา: จงหาจำนวนบุชเชลที่จะขนส่งจากแต่ละสวนไปยังแต่ละโรงงานแปรรูปเพื่อให้จำนวน bushel-mile ทั้งหมดที่ต้องขนส่งผลไม้ต่ำที่สุด ว่าควรขนส่งจากไหนไปไหนในปริมาณเท่าไร?



เอกสารอ้างอิง

- Bowersox, D. J., Closs, D. J., Cooper, M. B., & Bowersox, J. C. (2019). *Supply chain logistics management* (5th ed.). McGraw-Hill Education.
- Cachon, G. P., & Fisher, M. (2000). Supply chain inventory management and the value of shared information. *Management Science*, 46(8), 1032–1048. <https://doi.org/10.1287/mnsc.46.8.1032.12029>
- Chopra, S., & Meindl, P. (2019). *Supply chain management: Strategy, planning, and operation* (7th ed.). Pearson.
- Farahani, R. Z., SteadieSeifi, M., & Asgari, N. (2010). Multiple criteria facility location problems: A survey. *Applied Mathematical Modelling*, 34(7), 1689–1709. <https://doi.org/10.1016/j.apm.2009.10.005>
- Govindan, K., & Soleimani, H. (2017). A review of reverse logistics and closed-loop supply chains: A journal of cleaner production focus. *Journal of Cleaner Production*, 142, 371–384. <https://doi.org/10.1016/j.jclepro.2016.03.126>.
- Grant, D. B., Trautrim, A., & Wong, C. Y. (2022). *Sustainable logistics and supply chain management: Principles and practices for sustainable operations and management* (3rd ed.). Kogan Page.
- Harrison, A., Skipworth, H., van Hoek, R., & Aitken, J. (2019). *Logistics management and strategy: Competing through the supply chain* (6th ed.). Pearson.
- Ivanov, D., & Dolgui, A. (2020). Viability of intertwined supply networks: Extending the supply chain resilience angles towards survivability. *International Journal of Production Research*, 58(10), 2904–2915. <https://doi.org/10.1080/00207543.2020.1750727>.
- Ivanov, D., & Dolgui, A. (2020). A digital supply chain twin for managing the disruption risks and resilience in the era of Industry 4.0. *Production Planning & Control*, 32(9), 775–788. <https://doi.org/10.1080/09537287.2020.1768450>.

เอกสารอ้างอิง (ต่อ)

- Ivanov, D., & Dolgui, A. (2021). OR-methods for coping with the ripple effect in supply chains during COVID-19 pandemic: Managerial insights and research implications. *International Journal of Production Economics*, 232, 107921.
- Ivanov, D. (2021). Supply chain viability and the COVID-19 pandemic: A conceptual and formal generalisation of four major adaptation strategies. *International Journal of Production Research*, 59(12), 3535–3552.
- Liu, K. Y. (2022). *Supply Chain Analytics: Concepts, Techniques and Applications*. (1st ed.) Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-92224-5>.
- Lukardi, C., & Hamsal, M. (2020). *Warehouse selection using center-of-gravity method in minimizing transportation cost*. In *Contemporary Research on Business and Management*. Routledge / Taylor & Francis.
- Mentzer, J. T., Min, S., & Bobbitt, L. M. (2004). Toward a unified theory of logistics. *International Journal of Physical Distribution & Logistics Management*, 34(8), 606–627. <https://doi.org/10.1108/09600030410557758>
- Novack, R. A., Gibson, B. J., Suzuki, Y., & Coyle, J. J. (2019). *Transportation: A global supply chain perspective* (9th ed.). Cengage.
- Oeser, G., & Romano, P. (2016). *An empirical examination of the assumptions of the Square Root Law for inventory centralisation and decentralisation*. *International Journal of Production Research*, 54(8), 2298-2319. <https://doi.org/10.1080/00207543.2015.1071895>.
- Rodrigue, J.-P. (2020). *The geography of transport systems* (5th ed.). Routledge.
- Rushton, A., Croucher, P., & Baker, P. (2022). *The handbook of logistics and distribution management* (7th ed.). Kogan Page.

เอกสารอ้างอิง (ต่อ)

Şahin, M. (2021). *Location selection by multi-criteria decision-making methods based on objective and subjective weightings*. *Knowledge and Information Systems*, 63(8), 1991-2021. <https://doi.org/10.1007/s10115-021-01588-y>.

Simchi-Levi, D., Kaminsky, P., & Simchi-Levi, E. (2021). *Designing and managing the supply chain: Concepts, strategies, and case studies* (4th ed.). McGraw-Hill.

Singh, R. K., Chaudhary, N., & Saxena, N. (2018). Selection of warehouse location for a global supply chain: A case study. *IIMB Management Review*, 30(4), 343-354. <https://doi.org/10.1016/j.iimb.2018.08.009>.

บทที่ 8

การวิเคราะห์ข้อมูลโลจิสติกส์เพื่อจัดเส้นทางขนส่ง Logistics Data Analytics for Vehicle Routing

หัวข้อ

- 8.1 คุณลักษณะของปัญหาการจัดเส้นทางขนส่ง
- 8.2 การจัดเส้นทางขนส่ง
- 8.3 การใช้โปรแกรมเชิงเส้นตรง (Linear Programming, LP) แก้ปัญหาการจัดเส้นทางขนส่ง
- 8.4 การใช้วิธีฮิวริสติก (Heuristic) แก้ปัญหาการจัดเส้นทางขนส่ง
- 8.5 การจัดเส้นทางขนส่งด้วยการเรียนรู้ของเครื่อง

การจัดเส้นทางขนส่ง (vehicle routing problem, VRP) เป็นหัวใจสำคัญของธุรกิจโลจิสติกส์ที่ต้องการความได้เปรียบในการแข่งขัน เพราะการจัดเส้นทางที่มีประสิทธิภาพจะช่วยให้ธุรกิจสามารถ (1) ลดต้นทุนอย่างมีนัยสำคัญ เนื่องจากการลดระยะทางและจำนวนเที่ยวรถช่วยประหยัดค่าใช้จ่ายด้านเชื้อเพลิง บำรุงรักษา และค่าแรงได้อย่างมาก ตัวอย่างเช่น บริษัทขนส่งพัสดุรายใหญ่แห่งหนึ่งสามารถลดต้นทุนเชื้อเพลิงได้ถึง 20% เพียงแค่ปรับปรุงเส้นทางส่งพัสดุ (2) เพิ่มความเร็วและความแม่นยำในการส่งมอบ เพราะสินค้าถึงมือลูกค้าตรงเวลาและครบถ้วน ทำให้ลูกค้าพึงพอใจและกลับมาใช้บริการซ้ำอีก เช่น บริษัทจัดส่งอาหารสามารถลดเวลาในการส่งอาหารจาก 45 นาที เหลือเพียง 30 นาที ทำให้ลูกค้าได้รับอาหารที่ร้อนและสดใหม่ (3) เพิ่มประสิทธิภาพการใช้ทรัพยากร เนื่องจากใช้รถขนส่งน้อยลง ลดการปล่อยมลพิษ และลดความเสียหายต่อสิ่งแวดล้อมเช่น บริษัทขนส่งขยะสามารถลดจำนวนรถบรรทุกที่วิ่งบนท้องถนนได้ถึง 30% ช่วยลดการจราจรติดขัดและมลพิษทางอากาศ (4) ปรับตัวเข้ากับสถานะที่เปลี่ยนแปลงได้อย่างรวดเร็ว เมื่อเกิดเหตุการณ์ไม่คาดคิด เช่น อุบัติเหตุหรือการจราจรติดขัด ระบบจัดเส้นทางที่ทันสมัยสามารถปรับเปลี่ยนเส้นทางได้ทันที ทำให้การส่งมอบสินค้าไม่ล่าช้า ตัวอย่างเช่น บริษัทขนส่งสินค้าสดสามารถหลีกเลี่ยงเส้นทางที่มี

การก่อสร้าง เพื่อให้สินค้าถึงปลายทางในสภาพที่สดใหม่ และ (6) เพิ่มขีดความสามารถในการแข่งขัน เพราะธุรกิจที่มีระบบการจัดเส้นทางที่ดี จะสามารถให้บริการลูกค้าได้อย่างรวดเร็ว แม่นยำ และตรงตามความต้องการ ซึ่งเป็นปัจจัยสำคัญในการสร้างความแตกต่างและเอาชนะคู่แข่ง (Braekers et. al., 2016)

อย่างไรก็ตาม การจัดเส้นทางขนส่งเป็นปัญหาที่ซับซ้อนอย่างยิ่ง เนื่องจากมีปัจจัยหลายอย่างที่ต้องพิจารณาพร้อมกัน เช่น จำนวนจุดส่ง ระยะทาง เวลาที่กำหนด และข้อจำกัดต่างๆ ทำให้จำนวนวิธีในการจัดเส้นทางเพิ่มขึ้นอย่างรวดเร็วตามจำนวนจุดส่งที่มากขึ้น การหาคำตอบที่ดีที่สุดในเวลาอันสั้นจึงเป็นเรื่องยากมาก หรือกล่าวได้ว่าเป็นปัญหาที่อยู่ในกลุ่ม NP-hard ซึ่งหมายความว่ายังไม่มีอัลกอริทึมใดที่สามารถหาคำตอบที่แน่นอนได้ภายในเวลาที่เหมาะสมสำหรับปัญหาขนาดใหญ่ ดังนั้นเพื่อแก้ไขปัญหา นักวิจัยได้คิดค้นเทคนิคฮิวริสติก ซึ่งเป็นวิธีการที่ช่วยให้เราค้นหาคำตอบที่เกือบสมบูรณ์ (near optimal solution) ได้ในเวลาอันสั้น แม้ว่าจะไม่สามารถรับประกันได้ว่าจะเป็นคำตอบที่ดีที่สุดเสมอไปก็ตาม วิธีฮิวริสติกทำงานโดยการกำหนดกฎเกณฑ์หรือเงื่อนไขบางอย่างเพื่อลดจำนวนตัวเลือกในการค้นหา ทำให้สามารถหาคำตอบได้เร็วขึ้น นอกจากนี้ในปัจจุบัน เทคโนโลยีปัญญาประดิษฐ์ (AI) และ การเรียนรู้ของเครื่อง (ML) ได้เข้ามามีบทบาทสำคัญในการพัฒนาวิธีการแก้ปัญหา VRP ให้มีประสิทธิภาพมากยิ่งขึ้น โดย AI สามารถเรียนรู้จากข้อมูลจำนวนมากและค้นพบรูปแบบที่ซับซ้อน เพื่อนำมาใช้ในการตัดสินใจในการจัดเส้นทางได้อย่างอัจฉริยะ (Bai et al., 2021)

8.1 ลักษณะเฉพาะของปัญหาการวางแผนเส้นทางขนส่ง

ลองจินตนาการว่าคุณเป็นผู้จัดการฝ่ายโลจิสติกส์ของบริษัทขนส่งสินค้าขนาดใหญ่ มีลูกค้ากระจายอยู่ทั่วประเทศ และคุณต้องการส่งสินค้าให้ถึงมือลูกค้าทุกคนอย่างรวดเร็วและประหยัดต้นทุนที่สุด ปัญหาที่คุณกำลังเผชิญอยู่นี้คือ ปัญหาการจัดเส้นทางขนส่งซึ่งไม่ใช่เพียงแค่อุปสรรคทางคณิตศาสตร์ที่ซับซ้อน แต่ยังเป็นหัวใจสำคัญของธุรกิจโลจิสติกส์ (Toth & Vigo, 2014) การแก้ปัญหา VRP อย่างมีประสิทธิภาพจะช่วยให้ธุรกิจสามารถลดต้นทุนการขนส่งได้อย่างมาก ตัวอย่างเช่น บริษัทขนส่งพัสดุรายใหญ่แห่งหนึ่งสามารถลดต้นทุนเชื้อเพลิงได้ถึง 20% เพียงแค่ปรับปรุงเส้นทางส่งพัสดุให้เหมาะสม ลักษณะเฉพาะของปัญหาการวางแผนเส้นทางขนส่ง (VRP) อธิบายโดยสรุปดังนี้

1. การค้นหาเส้นทางที่เหมาะสมที่สุด (Dantzig & Ramser, 1959) การจัดเส้นทางขนส่งให้มีประสิทธิภาพสูงสุดเปรียบเสมือนการแก้ปริศนาที่ซับซ้อน เพราะต้องคำนึงถึงปัจจัยหลาย

อย่างพร้อมกัน เช่น ระยะทางและเวลาที่ใช้ในการเดินทาง ต้นทุนในการดำเนินการขนส่ง ข้อจำกัดของยานพาหนะ และความต้องการของลูกค้า เป้าหมายหลักคือการ ลดต้นทุนโดยรวม ซึ่งรวมถึงค่าเชื้อเพลิง ค่าแรงงาน และเวลาที่สูญเสียไปกับการเดินทางที่ไม่มีประสิทธิภาพ การกำหนดเส้นทางที่ไม่เหมาะสมอาจนำไปสู่ความล่าช้าในการส่งมอบสินค้า ความเสียหายระหว่างการขนส่ง และอาจสูญเสียลูกค้าในระยะยาว

2. **ข้อจำกัดของทรัพยากร** ปัจจัยสำคัญที่บีบให้การจัดเส้นทางมีความซับซ้อนยิ่งขึ้น ทำให้การวางแผนเส้นทางขนส่งให้มีประสิทธิภาพสูงสุดเป็นเรื่องที่ซับซ้อนและท้าทาย เพราะนอกจากจะต้องคำนึงถึงระยะทางและเวลาในการเดินทางแล้ว นอกจากนี้ยังต้องเผชิญกับข้อจำกัดด้านทรัพยากรที่หลากหลาย ซึ่งมีผลกระทบโดยตรงต่อกระบวนการวางแผนเส้นทาง เช่น ความจุของยานพาหนะเพราะรถบรรทุกแต่ละคันมีขนาดและความจุในการบรรทุกที่ไม่เหมือนกัน หากสินค้ามีปริมาณมากหรือมีขนาดใหญ่ อาจจำเป็นต้องใช้รถบรรทุกหลายคันเพื่อรองรับการขนส่งสินค้า ซึ่งอาจส่งผลให้ต้นทุนการขนส่งสูงขึ้น และยังเพิ่มความเสี่ยงในการเกิดความล่าช้าในการจัดส่งสินค้า นอกจากนี้ยังต้องคำนึงถึงน้ำหนักบรรทุกสูงสุดตามกฎหมายกำหนดให้รถบรรทุกต้องบรรทุกน้ำหนักไม่เกินค่าที่ระบุไว้ เพื่อความปลอดภัยและปฏิบัติตามข้อบังคับทางกฎหมาย หากสินค้าน้ำหนักมากเกินไป อาจต้องแบ่งส่งเป็นหลายเที่ยว หรือใช้รถบรรทุกขนาดใหญ่ขึ้น เวลาในการให้บริการเนื่องจากพนักงานขับรถมีเวลาทำงานที่จำกัด และกฎหมายกำหนดเวลาพักรถเพื่อความปลอดภัย หากเส้นทางยาวเกินไป อาจทำให้พนักงานขับรถเหนื่อยล้าและเกิดอุบัติเหตุได้ นอกจากนี้ เวลาในการขนส่งยังอาจถูกจำกัดด้วยปัจจัยอื่นๆ เช่น ช่วงเวลาที่ลูกค้าสามารถรับสินค้าได้ (time windows) อีกข้อจำกัดหนึ่งก็คือชนิดของสินค้าเพราะสินค้าบางชนิดต้องการอุณหภูมิที่ควบคุม หรือต้องบรรจุหีบห่อเป็นพิเศษ ซึ่งจะจำกัดประเภทของยานพาหนะที่สามารถใช้ขนส่งได้ ปัจจัยสภาพเส้นทางหรือสภาพเส้นทางที่ไม่ดี เช่น ถนนขรุขระ หรือการจราจรที่หนาแน่น อาจทำให้ความเร็วในการเดินทางลดลง และเพิ่มโอกาสที่จะเกิดความเสียหายต่อสินค้า ปัจจัยสุดท้ายคือกฎหมายและข้อบังคับเพราะในการขนส่งจะต้องปฏิบัติตามกฎหมายและข้อบังคับ เช่น ข้อจำกัดเวลาในการขับขี ข้อห้ามผ่านบางเส้นทาง และข้อกำหนดด้านความปลอดภัย จะมีผลต่อการวางแผนเส้นทาง

3. **ประเภทของปัญหา VRP** มีความซับซ้อนและหลากหลายรูปแบบ โดยแต่ละรูปแบบจะมีข้อจำกัดและความซับซ้อนที่แตกต่างกันไป เพื่อให้เข้าใจง่ายขึ้น ลองนึกภาพสถานการณ์ต่อไปนี้ (1) CVRP (Capacitated VRP) เหมือนกับการส่งพัสดุให้กับลูกค้าหลายราย โดยรถบรรทุกแต่ละคันมีขนาดที่กำหนด หากพัสดุมีขนาดใหญ่หรือมีจำนวนมาก อาจจำเป็นต้องใช้รถบรรทุกหลายคันในการขนส่งสินค้า (2) VRPTW (VRP with Time Windows) (Bräysy & Gendreau, 2005) คล้ายกับการส่งอาหารให้กับลูกค้าแต่ละราย โดยลูกค้าแต่ละคนจะมีช่วงเวลาที่กำหนดในการรับอาหาร เช่น ลูกค้า A ต้องการรับอาหารระหว่าง 12:00-13:00 น. รถส่งอาหารต้องวางแผนเส้นทางให้ทันตามเวลาที่กำหนด (3) MDVRP (Multi-Depot VRP) เหมือนกับบริษัทขนส่งที่มีหลายสาขา โดยรถบรรทุกจะออกจากสาขาที่ใกล้กับลูกค้าที่สุดเพื่อให้สามารถลดระยะทางในการเดินทางและลดต้นทุนการขนส่งโดยรวม (4) VRPPD (VRP with Pick-Up and Delivery) คล้ายกับการขนส่งสินค้าระหว่างโรงงานและคลังสินค้า โดยรถบรรทุกต้องไปรับสินค้าที่โรงงานก่อน แล้วนำไปส่งที่คลังสินค้า (5) VRPMT (VRP with Multiple Trips) รถบรรทุกสามารถกลับไปที่คลังสินค้าเพื่อบรรจุสินค้าเพิ่มเติมได้หลายครั้ง (6) VRPSD (VRP with Stochastic Demands) ปริมาณสินค้าที่ลูกค้าต้องการอาจไม่แน่นอน ต้องมีการวางแผนเส้นทางที่ยืดหยุ่น (7) VRPHF (VRP with Heterogeneous Fleet) มีการใช้รถบรรทุกหลายประเภทที่มีขนาดและความจุแตกต่างกัน
4. **วิธีการแก้ปัญหา VRP** ต้องการวิธีการที่หลากหลายในการแก้ไข เนื่องจากปัญหา VRP เป็นปัญหาที่ซับซ้อนสูง โดยแบ่งออกเป็น 3 ประเภทหลักคือ (1) วิธีการหาคำตอบที่แน่นอน exact methods) วิธีการเหล่านี้มุ่งเน้นการค้นหาคำตอบที่ดีที่สุดอย่างแม่นยำ เช่น การใช้วิธีการโปรแกรมเชิงเส้น การแบ่งสาขาและลิมิต (branch and bound) และการโปรแกรมเชิงพลวัต (dynamic programming) แต่ข้อจำกัดคือ ใช้เวลานานในการคำนวณเมื่อเจอปัญหาที่มีขนาดใหญ่ (2) อัลกอริทึมเชิงฮิวริสติก (heuristic algorithms) เป็นวิธีการที่เน้นการหาคำตอบที่น่าพอใจหรือใกล้เคียงกับคำตอบที่ดีที่สุดในเวลาที่ยอมรับได้ โดยอาศัยกฎเกณฑ์และประสบการณ์ในการตัดสินใจ ตัวอย่างเช่น Clarke-Wright savings algorithm เป็นวิธีการที่ถูกใช้อย่างแพร่หลายในกระบวนการแก้ปัญหา VRP โดยจะคำนวณค่าประหยัดจากการรวมเส้นทางของลูกค้าที่อยู่ใกล้กันเข้าด้วยกัน (3) เมตาฮิวริสติก (metaheuristic algorithms) (Gendreau et al., 2010) เป็นวิธีที่ใช้เพื่อเพิ่มประสิทธิภาพและปรับปรุง

ผลลัพธ์ที่ได้จากวิธีการเชิงฮิวริสติกให้ดียิ่งขึ้น โดยเลียนแบบพฤติกรรมตามธรรมชาติ ตัวอย่างเช่น genetic algorithms (GA) ที่เลียนแบบกระบวนการทางพันธุกรรมในการแสวงหาคำตอบที่เหมาะสมที่สุด simulated annealing (SA) ที่เลียนแบบกระบวนการหลอมโลหะในการค้นหาคำตอบ ant colony optimization (ACO) ที่เลียนแบบพฤติกรรมของมดในการค้นหาอาหาร หรือ particle swarm optimization (PSO) ที่เลียนแบบพฤติกรรมของฝูงนกในการค้นหาอาหาร

5. ความท้าทายเกิดขึ้นเมื่อปัญหา VRP จะมีความซับซ้อนเพิ่มมากขึ้นเมื่อจำนวนลูกค้าและข้อจำกัดต่าง ๆ เพิ่มขึ้น ลองจินตนาการว่าเรามีลูกค้า 10 รายที่ต้องส่งสินค้าไปให้ และเราต้องการหาเส้นทางที่มีประสิทธิภาพที่สุดในการส่งสินค้าให้กับลูกค้าทั้ง 10 คน โดยเส้นทางนี้จะต้องเริ่มต้นและสิ้นสุดที่โกดังสินค้า (depot) การหาคำตอบที่ถูกต้องเปรียบเสมือนการต่อจิ๊กซอว์ที่มีชิ้นส่วนจำนวนมากและมีวิธีการต่อที่เป็นไปได้นับไม่ถ้วน (Laporte, 2009)
6. ทำไมปัญหา VRP ถึงซับซ้อนขนาดนั้น? เพราะปัจจัย 2 ประการ ได้แก่ ประการแรกจำนวนวิธีการเดินทางที่เป็นไปได้ เมื่อจำนวนลูกค้าเพิ่มขึ้น วิธีการจัดเส้นทางที่เป็นไปได้อีกจะยิ่งเพิ่มขึ้นตาม (Lenstra & Rinnooy Kan, 1981) เช่น ถ้ามีลูกค้า 5 ราย จะมีวิธีการจัดเส้นทางได้ถึง $5!$ ($5 \times 4 \times 3 \times 2 \times 1$) หรือ 120 วิธี หรือ $5!/2 = 60$ เส้นทางที่มีลำดับการเดินทางแตกต่างกันโดยไม่ซ้ำกัน (เนื่องจากเส้นทางไปและกลับอาจเป็นเส้นทางเดียวกันแต่กลับทิศทางกัน) การคำนวณความซับซ้อนของปัญหาเมื่อจำนวนจุดหมายเพิ่มขึ้นแสดงดังตารางที่ 8.1 และตัวอย่างการกำหนดเส้นทางที่ไม่ซ้ำกันแสดงดังตัวอย่างที่ 8.1 ประการที่สอง ข้อจำกัดต่างๆที่มีหลากหลาย นอกเหนือจากจำนวนลูกค้า ยังมีปัจจัยจำกัดอื่น ๆ ที่ต้องคำนึงถึง เช่น เวลาในการส่งมอบ ความจุของรถบรรทุก และสภาพการจราจร ซึ่งทำให้ปัญหาซับซ้อนยิ่งขึ้นไปอีก

ตารางที่ 8.1 ความซับซ้อนของปัญหา VRP

จำนวน จุดหมาย	จำนวนเส้นทาง (ไม่ซ้ำ)
4	12
5	60
10	1,814,400
13	3,113,510,400
15	653,837,184,000
20	$1.21645100408832 \times 10^{18}$
50	$1.5207046600856688 \times 10^{64}$

ที่มา: จัดทำโดยผู้แต่ง

ตัวอย่าง 8.1 การกำหนดเส้นทางโดยไม่ให้ซ้ำกัน

สมมติว่าเรามีรถส่งของ 1 คัน และมีลูกค้า 4 ราย (A B C และ D) ที่ต้องการให้เราส่งสินค้าไปให้ โดยรถจะเริ่มต้นการเดินทางจากศูนย์กระจายสินค้า (depot) เพื่อทำการจัดส่งสินค้าให้กับลูกค้าทั้ง 4 ราย และกลับมายังศูนย์กระจายสินค้าอีกครั้ง การระบุเส้นทางที่เป็นไปได้ทั้งหมดโดยไม่ให้เส้นทางซ้ำกัน (ซึ่งหมายถึงเส้นทางที่กลับทิศทางกันจะนับเป็นเส้นทางเดียวกัน) สามารถทำได้ดังนี้:

1. การระบุขอบเขตของปัญหา

จุดเริ่มต้นที่สำคัญที่สุดในการแก้ปัญหา VRP จะต้องระบุจำนวนลูกค้าและตำแหน่งของลูกค้า แต่ละรายอย่างชัดเจน เราจะนำข้อมูลมาประมวลผลเพื่อสร้างภาพรวมของพื้นที่การจัดส่งเพื่อให้ได้ผลลัพธ์ในการวางแผนเส้นทางที่แม่นยำ ยกตัวอย่างเช่น ถ้าต้องการจัดส่งสินค้าให้กับลูกค้าจำนวน 4 ราย ได้แก่ ลูกค้า A B C และ D โดยมีศูนย์กระจายสินค้า (depot) อยู่ที่จุด O การระบุจำนวนและที่ตั้งของลูกค้าแต่ละราย เพื่อให้เห็นภาพการกระจายตัวของลูกค้าบนแผนที่รวมถึงศูนย์กระจายสินค้าทำให้สามารถวางแผนเส้นทางได้อย่างมีประสิทธิภาพ

2. การคำนวณจำนวนเส้นทางทั้งหมด

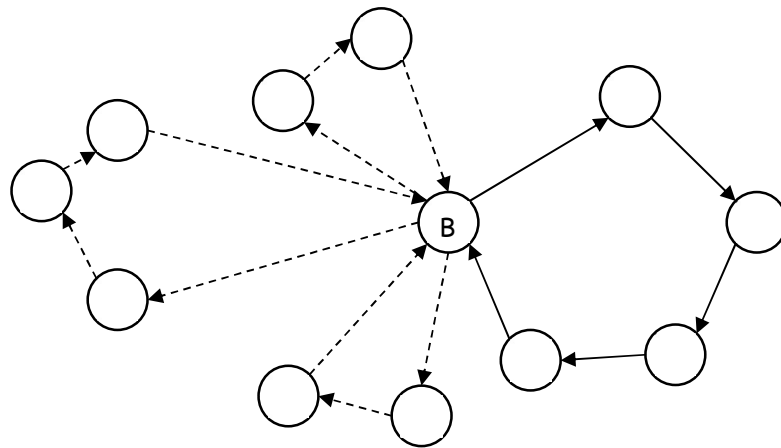
ถ้าเรามีลูกค้า 4 ราย (A B C และ D) ที่ต้องไปส่งของ เราจะไปส่งลูกค้าตามลำดับใดก็ได้ เช่น (A-B-C-D) (B-A-D-C) หรือ (C-D-A-B) เป็นต้น การเรียงลำดับที่แตกต่างกันก็จะได้เส้นทางที่แตกต่างกันไป (1) จำนวนลำดับทั้งหมด หากมีลูกค้า 4 ราย เราสามารถเรียงลำดับการส่งของให้กับลูกค้าได้ทั้งหมด $4!$ หรือ $4 \times 3 \times 2 \times 1 = 24$ ลำดับ (2) ตัดเส้นทางที่ซ้ำ ถ้าเราพิจารณาเส้นทาง A-B-C-D กับเส้นทาง D-C-B-A จะเห็นว่าเป็นเส้นทางเดียวกัน เพียงแค่กลับทิศทางการเดินทาง ดังนั้นเพื่อไม่ให้เกิดการนับซ้ำ เราจึงนำจำนวนลำดับทั้งหมดมาหารด้วย 2 (3) จำนวนเส้นทางที่แตกต่างกันทั้งหมดสำหรับลูกค้า 4 ราย คือ $24 / 2 = 12$ เส้นทาง

3. กำหนดเส้นทางที่ไม่ซ้ำกัน รายการของเส้นทางที่ไม่ซ้ำกันทั้งหมด (เส้นทางที่กลับทิศทางการเดินทางเป็นเส้นทางเดียวกัน) ได้แก่

1. O -> A -> B -> C -> D -> O
2. O -> A -> B -> D -> C -> O
3. O -> A -> C -> B -> D -> O
4. O -> A -> C -> D -> B -> O
5. O -> A -> D -> B -> C -> O
6. O -> A -> D -> C -> B -> O
7. O -> B -> C -> D -> A -> O
8. O -> B -> D -> C -> A -> O
9. O -> B -> A -> C -> D -> O
10. O -> B -> A -> D -> C -> O
11. O -> C -> D -> A -> B -> O
12. O -> C -> B -> A -> D -> O

8.2 การจัดเส้นทางขนส่ง (Vehicle Routing)

กำหนดให้มี n โหนดที่ต้องไปส่งของ แต่ละโหนดมีความต้องการ v_i ($i=1,2,...,n$) หน่วยของสินค้า เราต้องการหาเส้นทางที่เหมาะสมที่สุดสำหรับยานพาหนะในการเดินทางจากสถานี B ไปส่งสินค้ายัง n โหนด โดยคำนึงถึงปริมาณสินค้าที่ต้องส่งให้แต่ละโหนด และความจุของยานพาหนะ โดยที่ยานพาหนะทุกคันต้องเริ่มต้นและสิ้นสุดการเดินทางที่สถานี B เสมอ โดยที่ความจุของรถ V มากกว่าจำนวนสินค้าที่ต้องขนส่ง v_i แสดงดังภาพที่ 8.1



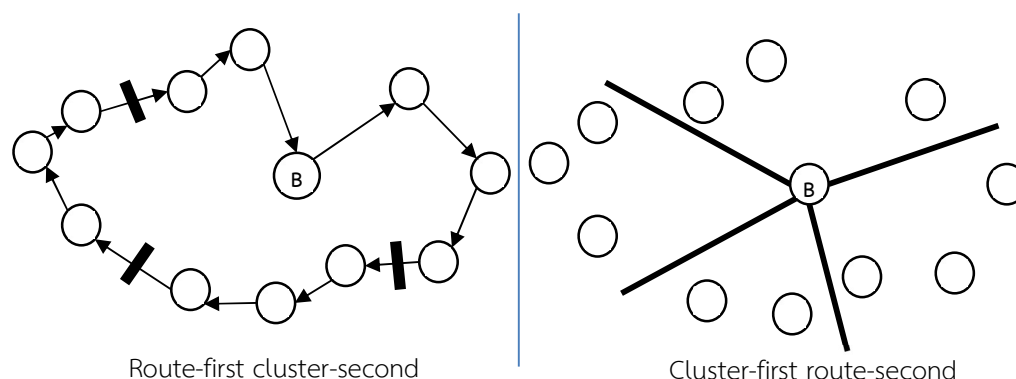
ภาพที่ 8.1 ตัวอย่างเส้นทางขนส่งอย่างง่าย

ที่มา: ภาพโดยผู้แต่ง

ปัญหา VRP (vehicle routing problem) หรือปัญหาการจัดเส้นทางยานพาหนะ นั้นมีวิธีการแก้ปัญหาที่หลากหลายถูกพัฒนาขึ้นอย่างต่อเนื่องในช่วงหลายทศวรรษที่ผ่านมา (Gillett & Miller, 1974; Fisher & Jaikumar, 1981; Solomon, 1987; Laporte & Semet, 2002)

โดยหลักแล้ววิธีการแก้ปัญหา VRP เหล่านี้สามารถแบ่งออกเป็น 2 กลุ่มหลักแสดงดังภาพที่ 8.2 ได้แก่ (1) จัดกลุ่มก่อน จัดเส้นทางทีหลัง โดยกลุ่มวิธีการนี้จะทำการแบ่งลูกค้าหรือจุดหมายปลายทางออกเป็นกลุ่มๆ ก่อน แล้วจึงทำการวางแผนเส้นทางสำหรับแต่ละกลุ่ม โดยพิจารณาปัจจัยต่างๆ เช่น ความจุของยานพาหนะ ระยะทาง และข้อจำกัดอื่นๆ (2) จัดเส้นทางแล้วจึงจัดกลุ่ม กลุ่มวิธีการนี้จะทำการวางแผนเส้นทางทั้งหมดก่อน แล้วจึงทำการแบ่งเส้นทางที่ได้ออกเป็นกลุ่มๆ โดย

พิจารณาจากความจุของยานพาหนะและข้อจำกัดอื่นๆ การเลือกใช้วิธีการใดขึ้นอยู่กับลักษณะของปัญหา VRP ที่ต้องการแก้ไข เช่น จำนวนลูกค้า ความซับซ้อนของเส้นทาง ข้อจำกัดด้านเวลา และทรัพยากรที่มีอยู่



ภาพที่ 8.2 จัดเส้นทางแล้วจึงจัดกลุ่ม; จัดกลุ่มแล้วจึงจัดเส้นทาง

ที่มา: ภาพโดยผู้แต่ง

8.3 การใช้โปรแกรมเชิงเส้นตรง (Linear Programming, LP) แก้ปัญหาการขนส่ง

LP เป็นเครื่องมืออันทรงพลังในการสร้างแบบจำลองทางคณิตศาสตร์ที่ชัดเจนและมีระบบ (Hillier & Lieberman, 2021) โดยเฉพาะอย่างยิ่งสำหรับปัญหาที่มีข้อจำกัดเป็นเชิงเส้น LP ช่วยให้เราสามารถนิยาม ฟังก์ชันวัตถุประสงค์ (เช่น การลดค่าใช้จ่ายหรือการเพิ่มผลกำไร) และ ข้อจำกัดต่างๆ ภายในระบบในรูปแบบของสมการทางคณิตศาสตร์ ทำให้เราสามารถวิเคราะห์ปัญหาได้อย่างเป็นระบบและตรวจสอบความถูกต้องของคำตอบได้ง่ายยิ่งขึ้น จุดเด่นที่สำคัญของ LP คือความสามารถในการหา "คำตอบที่เหมาะสมที่สุด" (optimal solution) ได้อย่างแม่นยำ เมื่อเรากำหนดเป้าหมายที่ชัดเจน เช่น การลดค่าใช้จ่ายในการขนส่งให้เหลือน้อยที่สุด หรือการลดจำนวนยานพาหนะที่ใช้ LP จะช่วยให้เราค้นหาแผนการดำเนินงานที่ดีที่สุดที่สอดคล้องกับเงื่อนไขที่กำหนดไว้ สำหรับปัญหา VRP ขนาดเล็กถึงกลาง ซึ่งมีจำนวนยานพาหนะและลูกค้าไม่มากเกินไป LP สามารถให้ผลลัพธ์ที่น่าพอใจและนำไปใช้งานได้จริงในเวลาที่เหมาะสม เมื่อเทียบกับวิธีการอื่น ๆ อย่างไรก็ตาม เมื่อขนาดของ

ปัญหา VRP เพิ่มขึ้น เช่น มีจำนวนลูกค้าและยานพาหนะเป็นจำนวนมาก LP อาจประสบปัญหาในการหาคำตอบที่ดีที่สุดภายในระยะเวลาที่กำหนด เนื่องจากความซับซ้อนของปัญหาที่เพิ่มขึ้น หนึ่งในปัญหา VRP ที่เป็นที่ยู้งานกันดีและถูกนำมาศึกษาอย่างกว้างขวางคือ ปัญหาการเดินทางของพนักงานขาย (traveling salesman problem, TSP) ซึ่งเป็นปัญหาคาสสิกที่เกี่ยวข้องกับการหาเส้นทางที่สั้นที่สุดในการเดินทางไปยังทุกเมืองและกลับมายังจุดเริ่มต้น เราจะมาทำความเข้าใจปัญหา TSP อย่างละเอียดในหัวข้อถัดไป

8.3.1 ปัญหาการเดินทางของพนักงานขาย (TSP)

ปัญหาการเดินทางของพนักงานขาย (TSP) เป็นปัญหาคาสสิกที่ทำให้นักวิทยาศาสตร์และวิศวกรมาหลายชั่วอายุคน (Dantzig et al., 1954) เปรียบเสมือนปริศนาที่ซับซ้อนที่รอคอยการไขปริศนา โดยในปัญหา TSP นั้น พนักงานขายต้องออกเดินทางจากสำนักงานเพื่อไปเยี่ยมลูกค้าที่กระจายอยู่ในหลายพื้นที่ โดยมีเงื่อนไขว่าต้องไปเยี่ยมลูกค้าทุกคนเพียงครั้งเดียว และกลับมายังสำนักงานในที่สุด เป้าหมายสูงสุดคือการค้นหาเส้นทางที่สั้นที่สุด เพื่อประหยัดทั้งเวลาและทรัพยากรในการเดินทาง แม้ว่าปัญหา TSP จะฟังดูเรียบง่าย (Lawler et al., 1985) เพียงแค่ค้นหาเส้นทางที่สั้นที่สุดในการเยี่ยมชมทุกเมืองและกลับมายังจุดเริ่มต้น แต่ความซับซ้อนของปัญหานี้กลับซ่อนอยู่เบื้องหลังความเรียบง่ายนี้ การหาเส้นทางที่เหมาะสมที่สุดนั้นไม่ใช่เรื่องง่ายเลย โดยเฉพาะอย่างยิ่งเมื่อมีปัจจัยต่างๆ เข้ามาเกี่ยวข้อง เช่น ระยะทางที่แตกต่างกันในแต่ละทิศทาง ข้อจำกัดด้านเวลา หรือความจุของพาหนะเข้ามาเกี่ยวข้อง ทำให้ TSP กลายเป็นหนึ่งในปัญหาที่ท้าทายที่สุดในด้านการเพิ่มประสิทธิภาพและการวางแผนเส้นทาง ดังจะเห็นได้ว่าการขยายงานวิจัยที่เพิ่มตัวแปรและความซับซ้อนให้กับปัญหา TSP แบบดั้งเดิม ยกตัวอย่างเช่น TSP ที่ไม่สมมาตร (asymmetric TSP, ATSP) TSP ที่มีขอบเขตเวลา (TSP with time windows, TSPTW) และ TSP พนักงานขายหลายคน (multiple TSP, MTSP) เป็นต้น

นอกจากข้อจำกัดต่าง ๆ ดังที่กล่าวมาแล้วข้างต้น เมื่อจำนวนเมืองที่ต้องเยี่ยมชมเพิ่มขึ้น จำนวนเส้นทางที่เป็นไปได้ทั้งหมดจะเพิ่มขึ้นอย่างมหาศาล (Applegate et al., 2006) ทำให้การค้นหาเส้นทางที่ดีที่สุดนั้นเปรียบเสมือนการหาเข็มในมหาสมุทร การคำนวณเพื่อหาเส้นทางที่เหมาะสมที่สุดจึงใช้เวลานานมากขึ้นตามลำดับ จนอาจเป็นไปได้ในทางปฏิบัติหากมีจำนวนเมืองมากเกินไป เพื่อรับมือกับความซับซ้อนของปัญหา TSP นักวิจัยได้พัฒนาวิธีการต่าง ๆ มาใช้ในการค้นหาเส้นทางที่เหมาะสมที่สุด อาทิ การใช้เทคนิคทางคณิตศาสตร์อย่างการเขียนโปรแกรมเชิงเส้น

หรือการจำลองเหตุการณ์แบบ Monte Carlo นอกจากนี้ ยังมีการนำเอาแนวคิดจากธรรมชาติมาประยุกต์ใช้ เช่น อัลกอริทึมเชิงพันธุกรรมที่เลียนแบบกระบวนการวิวัฒนาการ อัลกอริทึมอาณานิคมมดที่เลียนแบบพฤติกรรมของมดในการค้นหาอาหาร หรืออัลกอริทึมค้างคาวที่เลียนแบบการใช้เสียงของค้างคาวในการนำทาง ซึ่งวิธีการเหล่านี้ล้วนมีจุดเด่นและข้อจำกัดที่แตกต่างกันไป (Osaba et al. 2020)

หนึ่งในรูปแบบที่โดดเด่นสำหรับการแก้ปัญหา TSP คือการกำหนดรูปแบบของ Miller-Tucker-Zemlin (MTZ) (Miller et al. 1960) ที่อธิบายไว้ในปี 1960 สามารถแสดงได้ดังนี้

1. สมมติฐานเบื้องต้น (Hypothesis)

ในเครือข่ายที่มี n จุด พนักงานขายเริ่มการเดินทางที่จุด 1 (ซึ่งเป็นจุดศูนย์กลาง) และต้องเยือนแต่ละจุดเพียงครั้งเดียวแล้วกลับมายังจุด 1

2. ตัวแปร (Variables)

x_{ij} : ตัวแปรทวิภาค (binary variable) ที่แสดงว่ามีการเดินทางจากจุด i ไปยังจุด j หรือไม่ (1 ถ้ามีการเดินทาง, 0 ถ้าไม่มี)

u_i : ตัวแปรที่ช่วยป้องกันวงจรย่อยที่ไม่พึงประสงค์ (sub-tour) โดยแสดงลำดับของจุดที่เยือน

c_{ij} : ต้นทุนที่เกี่ยวข้อง (ตย. ระยะทาง) ของ (i,j)

3. ฟังก์ชันวัตถุประสงค์ (Objective Function)

เป้าหมายคือ ลดระยะทางของการเดินทางทั้งหมดให้น้อยที่สุด

$$\text{Minimize: } \sum_{i=0}^n \sum_{j=0}^n c_{ij} x_{ij} \quad (8.1)$$

โดยที่:

c_{ij} คือระยะทางหรือค่าขนส่งจากจุด i ไปยังจุด j

x_{ij} คือค่าที่เป็น 1 หรือ 0 ตามการเลือกเส้นทาง

4. ข้อจำกัด (Constraints)

ข้อจำกัดของการเยือนแต่ละจุดเพียงครั้งเดียว:

$$\sum_{j=1, j \neq i}^n x_{ij} = 1 \quad \forall i = 1, 2, \dots, n \quad (8.2)$$

$$\sum_{i=1, i \neq j}^n x_{ij} = 1 \quad \forall j = 1, 2, \dots, n \quad (8.3)$$

ข้อจำกัดของการป้องกันวงจรรย่อย (sub-tour elimination constraints):

$$u_i - u_j + nx_{ij} \leq n - 1 \quad \forall i \neq j, 2 \leq i, j \leq n \quad (8.4)$$

ข้อจำกัดของตัวแปร:

$$\begin{aligned} x_{ij} &\in \{0, 1\} \quad \forall i, j \\ u_i &\geq 0 \quad \forall i \end{aligned}$$

แม้ว่าการเขียนโปรแกรมเชิงเส้นจะเป็นวิธีการที่แม่นยำในการแก้ปัญหา TSP แต่ในโลกแห่งความเป็นจริงซึ่งมีข้อจำกัดด้านเวลาและทรัพยากร การหาคำตอบที่ดีที่สุดอาจไม่ใช่สิ่งที่จำเป็นเสมอไป อัลกอริทึมฮิวริสติกจึงเข้ามามีบทบาทสำคัญในการค้นหาคำตอบที่ "ดีพอ" ภายในระยะเวลาที่กำหนด ตัวอย่างเช่น ในอุตสาหกรรมโลจิสติกส์ การใช้ฮิวริสติกในการวางแผนเส้นทางขนส่งสามารถช่วยลดต้นทุนเชื้อเพลิงและเวลาในการเดินทางได้อย่างมีนัยสำคัญ

ตัวอย่าง 8.2 การแก้ปัญหา TSP ด้วย LP

สมมติว่ามี 4 เมือง โดยมีระยะทางระหว่างเมืองดังนี้:

	×	To	×	×	×
From	1	2	3	4	
1	0	13	21	26	
2	10	0	29	20	
3	30	20	0	5	
4	12	30	7	0	

ระยะทางจากเมือง i ไปเมือง j จะเขียนแทนด้วย c_{ij} และเราต้องการหาเส้นทางที่มีระยะทางรวมน้อยที่สุด

หมายเหตุ: ระยะทางขาไปจะไม่เท่ากับระยะทางขากลับ เช่น ระยะทางจากเมือง 1 ไป 2 เท่ากับ 13 หน่วย ในขณะที่ ระยะทางจากเมือง 2 ไป 1 จะเท่ากับ 10 หน่วย เป็นต้น

1. ตัวแปรการตัดสินใจ (Decision Variables)

กำหนดตัวแปร x_{ij} ซึ่งเป็นค่าศูนย์หรือหนึ่ง (binary variable) แทนการตัดสินใจว่าจะเดินทางจากเมือง i ไปเมือง j หรือไม่ โดย:

$$x_{ij} = \begin{cases} 1 & \text{หากเดินทางจากเมือง } i \text{ ไปยังเมือง } j \\ 0 & \text{หากไม่เดินทาง} \end{cases}$$

2. ฟังก์ชันวัตถุประสงค์ (Objective Function)

ฟังก์ชันวัตถุประสงค์คือ การลดระยะทางที่ใช้เดินทางทั้งหมดให้น้อยที่สุด:

$$\text{Minimize: } \sum_{i=1}^4 \sum_{j=1}^4 c_{ij} x_{ij}$$

3. ข้อจำกัด (Constraints)

ข้อจำกัดในการเข้าและออกจากแต่ละเมือง แต่ละเมืองต้องถูกเยี่ยมชมครั้งเดียว และต้องเดินทางออกจากเมืองนั้นครั้งเดียวเท่านั้น

$$\sum_{j=1}^4 x_{ij} = 1 \quad \forall i$$

$$\sum_{i=1}^4 x_{ij} = 1 \quad \forall j$$

ข้อจำกัดในการป้องกันการเดินทางเป็นวงย่อย (subtour elimination constraints)

เพื่อให้การเดินทางเป็นไปอย่างมีประสิทธิภาพและไม่เกิดการวนกลับซ้ำซ้อนในเส้นทาง เราจะใช้ตัวเลขพิเศษที่เรียกว่า "ลำดับที่ (u_i)" ให้กับแต่ละจุดหมาย (ยกเว้นจุดเริ่มต้น) ตัวเลขนี้จะบอกให้เราทราบว่าเราจะต้องเดินทางไปยังจุดหมายใดเป็นอันดับที่เท่าไรในแผนการเดินทางทั้งหมด การกำหนดลำดับที่นี้จะช่วยให้เราสามารถตรวจสอบได้ว่าเส้นทางที่วางแผนไว้นั้นเป็นไปตามเงื่อนไขที่กำหนดหรือไม่ และไม่มีการวนกลับไปยังจุดหมายที่เคยไปมาแล้ว

$$u_i - u_j + 4x_{ij} \leq 3 \quad \forall i \neq j, \quad 2 \leq i, j \leq 4$$

4. การตั้งโมเดลทั้งหมด

$$\text{Minimize } z = 13X_{12} + 21X_{13} + 26X_{14} + 10X_{21} + 29X_{23} + 20X_{24} + 30X_{31} + 20X_{32} + 5X_{34} + 12X_{41} + 30X_{42} + 7X_{43}$$

Subject to:

$$X_{12} + X_{13} + X_{14} = 1$$

$$X_{21} + X_{23} + X_{24} = 1$$

$$X_{31} + X_{32} + X_{34} = 1$$

$$X_{41} + X_{42} + X_{43} = 1$$

$$X_{21} + X_{31} + X_{41} = 1$$

$$X_{12} + X_{32} + X_{42} = 1$$

$$X_{13} + X_{23} + X_{43} = 1$$

$$X_{14} + X_{24} + X_{34} = 1$$

$$4X_{23} + U_2 - U_3 \leq 3$$

$$4X_{24} + U_2 - U_4 \leq 3$$

$$4X_{32} - U_2 + U_3 \leq 3$$

$$4X_{34} + U_3 - U_4 \leq 3$$

$$4X_{42} - U_2 + U_4 \leq 3$$

$$4X_{43} - U_3 + U_4 \leq 3$$

$$U_2, U_3, U_4 \in \mathbb{Z}^+$$

$$x_{ij} \in \{0, 1\}, \forall_{ij} \in [1, 4]$$

5. คำตอบที่ได้จากการแก้ปัญหา LP

ตารางที่ 8.2 คำตอบของตัวอย่างที่ 8.2

#	ตัวแปรตัดสินใจ	คำตอบ	ต้นทุน
1	X_{12}	1	13
2	X_{13}	0	21
3	X_{14}	0	26
4	X_{21}	0	10
5	X_{23}	1	29
6	X_{24}	0	20
7	X_{31}	0	30
8	X_{32}	0	20
9	X_{34}	1	5
10	X_{41}	1	12
11	X_{42}	0	30
12	X_{43}	0	7
13	U_2	0	0
14	U_3	1	0
15	U_4	2	0

ที่มา: จัดทำโดยผู้แต่ง

คำตอบที่ได้คือ: ระยะทางรวม 59 โดยมีเส้นทางดังนี้ 1 -> 2 -> 3 -> 4 -> 1

8.4 การใช้เทคนิคฮิวริสติก (Heuristic) แก้ปัญหาการขนส่ง

เทคนิคฮิวริสติก คือวิธีการแก้ปัญหาที่อาศัยประสบการณ์หรือกฎเกณฑ์ที่ใช้กันทั่วไป ทำให้เราสามารถหาคำตอบที่ใกล้เคียงกับคำตอบที่ดีที่สุดได้อย่างรวดเร็ว แม้จะไม่ได้แม่นยำ 100% แต่ก็มีประโยชน์ในการตัดสินใจและแก้ปัญหาในสถานการณ์ที่จำกัดเวลา

ตัวอย่างของเทคนิคฮิวริสติกที่ใช้ในการแก้ปัญหาการจัดเส้นทาง

Nearest Neighbor Heuristic เป็นวิธีที่เลือกจุดหมายปลายทางที่ใกล้ที่สุดจากตำแหน่งปัจจุบันเสมอ ทำให้การวางแผนเส้นทางเป็นไปอย่างรวดเร็ว

Insertion Heuristic วิธีนี้จะค่อยๆ เพิ่มจุดหมายปลายทางเข้าไปในเส้นทางที่มีอยู่ โดยเลือกตำแหน่งที่ทำให้ระยะทางรวมสั้นที่สุด

Clark-Wright Saving Algorithm วิธีนี้จะพิจารณาการรวมเส้นทางย่อยต่างๆ เข้าด้วยกันเพื่อลดระยะทางที่ต้องเดินทางทั้งหมด

เทคนิคฮิวริสติกเป็นเครื่องมือที่มีประโยชน์ในการแก้ปัญหาต่างๆ โดยเฉพาะอย่างยิ่งปัญหาที่ซับซ้อนและต้องการคำตอบที่รวดเร็ว แม้ว่าผลลัพธ์ที่ได้อาจไม่ใช่คำตอบที่ดีที่สุด แต่ก็เพียงพอสำหรับการใช้งานในหลายๆ กรณี

8.4.1 วิธีจัดเส้นทางแบบประหยัด (Clark-Wright's Saving Algorithm for the VRP)

วิธี **Clark-Wright Saving Method** เป็นหนึ่งในเทคนิคที่นิยมใช้ในการแก้ปัญหาการจัดเส้นทางยานพาหนะ เพื่อหาเส้นทางขนส่งที่สั้นที่สุดและประหยัดต้นทุนมากที่สุด โดยวิธีนี้จะทำงานผ่านขั้นตอนดังต่อไปนี้ (Clarke & Wright, 1964):

1. **คำนวณระยะทางเบื้องต้น** โดยคำนวณระยะทางจากจุดเริ่มต้นไปยังจุดหมายทุกแห่ง (สมมติว่าจุดเริ่มต้นคือศูนย์กลางการกระจายสินค้า) โดยใช้ระยะทางแบบยูคลิเดียน (Euclidean Distance) ดังสมการที่ 8.5

$$d(i, j) = \sqrt{(x_i - x_j)^2 + (y_i - y_j)^2} \quad (8.5)$$

2. **คำนวณค่า Saving (การประหยัด)** สำหรับแต่ละคู่ของจุดหมายด้วยสูตร:

$$S_{ij} = d_{0i} + d_{0j} - d_{ij} \quad (8.6)$$

โดย S_{ij} คือค่าการประหยัด, d_{0i} คือระยะทางจากจุดเริ่มต้นไปยังจุดหมาย i d_{0j} คือระยะทางจากจุดเริ่มต้นไปยังจุดหมาย j และ d_{ij} คือระยะทางระหว่างจุดหมาย i และ j

3. จัดเรียงค่าการประหยัด โดยหลังจากที่ได้ค่าประหยัดที่คำนวณจากสมการที่ 8.6 ทำการจัดเรียงจากค่ามากที่สุดไปน้อยที่สุด

4. รวมเส้นทาง โดยทำตามเงื่อนไขหลัก 3 เงื่อนไขดังต่อไปนี้

- 1) เริ่มจากคู่ของจุดหมายที่มีค่าการประหยัดสูงสุดรวมเข้าด้วยกันเป็นเส้นทางเดียว
- 2) ตรวจสอบเงื่อนไข เช่น ข้อจำกัดของความจุของยานพาหนะและข้อกำหนดของเส้นทาง
- 3) ทำซ้ำขั้นตอนนี้อันจนกว่าจะไม่มีค่าการประหยัดเพิ่มเติม

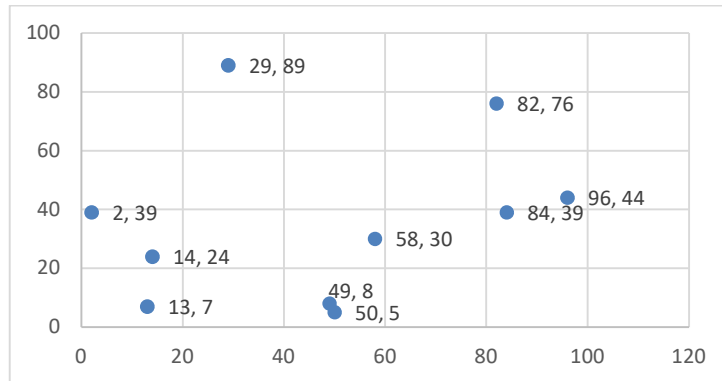
5. ตรวจสอบและปรับปรุง เริ่มจากตรวจสอบเส้นทางที่ได้เพื่อให้แน่ใจว่าปฏิบัติตามข้อกำหนดทั้งหมดและปรับปรุงเส้นทางหากจำเป็น

ตัวอย่าง 8.3 การคำนวณค่าการประหยัด

1. จุดเริ่มต้น 0 จุดหมาย A และ B
 - ระยะทาง $d_{0A}=10$, $d_{0B}=15$, $d_{AB}=12$
 - ค่าการประหยัด $S_{AB}=d_{0A}+d_{0B}-d_{AB}=10+15-12=13$
2. จัดเรียงค่าการประหยัดและรวมเส้นทาง:
 - เส้นทาง $0 \rightarrow A \rightarrow B \rightarrow 0$

วิธีนี้ช่วยลดระยะทางและต้นทุนในการขนส่งโดยการรวมเส้นทางที่มีประสิทธิภาพสูงสุด

ตัวอย่าง 8.4 การจัดเส้นทางด้วยวิธี Saving Algorithm



ภาพที่ 8.3 พิกัดตำแหน่งที่ตั้งของลูกค้าประกอบตารางที่ 8.3

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม MS Excel)

ลูกค้ามีตำแหน่งที่ตั้ง โหนด 1-9 รถบรรทุกแต่ละคันมีความจุเท่ากันที่ 50 หน่วย และ Depot อยู่ที่ตำแหน่ง 0 โดยมีพิกัดและความต้องการของแต่ละโหนดแสดงดังตารางต่อไปนี้

ตารางที่ 8.3 ตารางข้อมูลสำหรับตัวอย่างที่ 8.4

โหนด	x	y	ความต้องการ
0	82	76	0
1	96	44	19
2	50	5	21
3	49	8	6
4	13	7	19
5	29	89	7
6	58	30	12
7	84	39	16
8	14	24	6
9	2	39	16

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 1: ทำการคำนวณหาระยะทางเดินทางระหว่างโหนดใช้สมการที่ 8.5

ระยะทางจาก Depot (node 0) ไปยัง node 1 จะมีระยะทางเดียวกับ ระยะทางจาก node 1 ไป node 0 โดยสามารถคำนวณโดยใช้สมการที่ 8.5 ได้ดังนี้

$$\begin{aligned} d(0,1) &= \sqrt{(x_0 - x_1)^2 + (y_0 - y_1)^2} = \sqrt{(82 - 96)^2 + (76 - 44)^2} \\ &= \sqrt{(-14)^2 + (32)^2} \approx 34.93 \end{aligned}$$

ผลการคำนวณระยะทางทั้งหมดแสดงดังตารางที่ 8.4

ตารางที่ 8.4 ระยะทางระหว่างการเดินทางจาก โหนด i ไปยัง โหนด j

	0	1	2	3	4	5	6	7	8	9
0	0.00									
1	34.93	0.00								
2	77.88	60.31	0.00							
3	75.58	59.20	3.16	0.00						
4	97.58	90.87	37.05	36.01	0.00					
5	54.57	80.71	86.59	83.43	83.55	0.00				
6	51.88	40.50	26.25	23.77	50.54	65.74	0.00			
7	37.05	13.00	48.08	46.75	77.88	74.33	27.51	0.00		
8	85.60	84.40	40.71	38.48	17.03	66.71	44.41	71.59	0.00	
9	88.14	94.13	58.82	56.30	33.84	56.82	56.72	82.00	19.21	0.00

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 2: คำนวณค่าประหยัด(Saving) โดยสามารถคำนวณโดยใช้สมการที่ 8.6 ได้ดังนี้

$$S_{12} = d_{01} + d_{02} - d_{12} = 34.93 + 77.88 - 60.31 = 52.5$$

ผลประหยัดทั้งหมดระหว่างโหนดแสดงดังตารางที่ 8.5

ตารางที่ 8.5 ผลประหยัดระหว่างการเดินทางต่อเนื่องจาก โหนด i ไปยัง โหนด j

	1	2	3	4	5	6	7	8
2	52.50							
3	51.31	150.30						
4	41.64	138.40	137.15					
5	8.79	45.86	46.72	68.61				
6	46.32	103.51	103.70	98.93	40.71			
7	58.98	66.85	65.88	56.76	17.29	61.42		
8	36.13	122.78	122.70	166.16	73.47	93.08	51.07	
9	28.94	107.20	107.42	151.88	85.89	83.31	43.20	154.54

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 3: จัดเรียงค่าการประหยัดจากมากไปน้อย

ทำการจัดเรียงผลประหยัดจากมากไปน้อยได้ผลลัพธ์การจัดเรียงแสดงดังตารางที่ 8.6

ตารางที่ 8.6 จัดเรียงผลประหยัดระหว่างการเดินทางต่อเนื่องจาก โหนด i ไปยัง โหนด j

Rank	1	2	3	4	5	6	7	8	9	10	11	12	13
i	4	8	4	2	2	3	2	3	3	2	3	2	4
j	8	9	9	3	4	4	8	8	9	9	6	6	6
Saving	166	155	152	150	138	137	123	123	107	107	104	104	99

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 8.6 (ต่อ) จัดเรียงผลประหยัดระหว่างการเดินทางต่อเนื่องจาก โหนด i ไปยัง โหนด j

Rank	14	15	16	17	18	19	20	21	22	23	24	25	26
i	6	5	6	5	4	2	3	6	1	4	1	1	7
j	8	9	9	8	5	7	7	7	7	7	2	3	8
Saving	93	86	83	73	69	67	66	61	59	57	52	51	51

ที่มา: จัดทำโดยผู้แต่ง

ตารางที่ 8.6 (ต่อ) จัดเรียงผลประหยักระหว่างการเดินทางต่อเนื่องจาก โหนด i ไปยัง โหนด j

Rank	27	28	29	30	31	32	33	34	35	36
i	3	1	2	7	1	5	1	1	5	1
j	5	6	5	9	4	6	8	9	7	5
Saving	47	46	46	43	42	41	36	29	17	9

ที่มา: จัดทำโดยผู้แต่ง

ขั้นตอนที่ 4: การรวมเส้นทางที่ละขั้นตอน โดยเริ่มจากผลประหยัมากที่สุดที่สุดในตารางที่ 8.6

อันดับ	(i, j)	ความจุ (หน่วย)	เส้นทาง	หมายเหตุ
1	(4, 8)	$19+6 = 25$	0-4-8-0	
2	(8, 9)	$25+16 = 41$	0-4-8-9-0	
3	(4, 9)	-	-	(4, 9) อยู่ในเส้นทางแล้ว ตัดออก
4	(2, 3)	$21+6 = 27$	0-2-3-0	(2, 3) ไม่อยู่ในเส้นทางใดๆ สร้างเส้นทางใหม่
5	(2, 4)	-	-	รวมไม่ได้ เกินความจุของรถ
6	(3, 4)	-	-	รวมไม่ได้ เกินความจุของรถ
7	(2, 8)	-	-	รวมไม่ได้ 8 เป็น จุดภายใน
8	(3, 8)	-	-	รวมไม่ได้ 8 เป็น จุดภายใน
9	(3, 9)	-	-	รวมไม่ได้ เกินความจุของรถ
10	(2, 9)	-	-	รวมไม่ได้ เกินความจุของรถ
11	(3, 6)	$27+12 = 39$	0-2-3-6-0	
12	(2, 6)	-	-	(2, 6) อยู่ในเส้นทางแล้ว ตัดออก
13	(4, 6)	-	-	รวมไม่ได้ เกินความจุของรถ
14	(6, 8)	-	-	รวมไม่ได้ 8 เป็น จุดภายใน
15	(5, 9)	$41+7 = 48$	0-4-8-9-5-0	
16	(6, 9)	-	-	รวมไม่ได้ 9 เป็น จุดภายใน
17	(5, 8)	-	-	(5, 8) อยู่ในเส้นทางแล้ว ตัดออก
18	(4, 5)	-	-	(4, 5) อยู่ในเส้นทางแล้ว ตัดออก
19	(2, 7)	-	-	รวมไม่ได้ เกินความจุของรถ
20	(3, 7)	-	-	รวมไม่ได้ 3 เป็น จุดภายใน
21	(6, 7)	-	-	รวมไม่ได้ เกินความจุของรถ
22	(1, 7)	$19+16 = 35$	0-1-7-0	(1, 7) ไม่อยู่ในเส้นทางใดๆ สร้างเส้นทางใหม่

หยุดเพราะครบทุก node 1-9

ที่มา: จัดทำโดยผู้แต่ง

หลังจากการใช้ Saving Algorithm จะได้เส้นทางขนส่ง 3 เส้นทาง ดังนี้

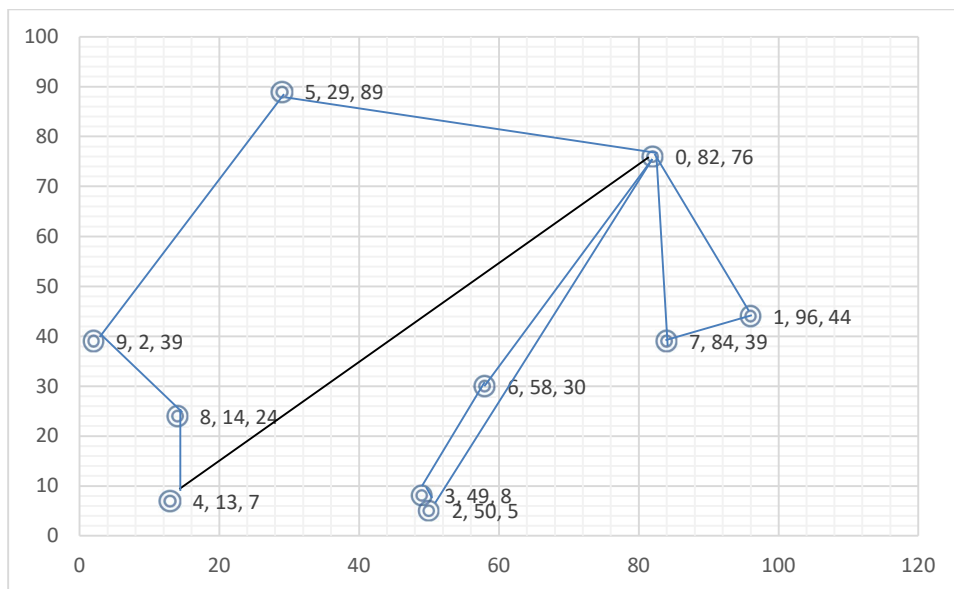
เส้นทางที่ 1 จาก $0 \rightarrow 4 \rightarrow 8 \rightarrow 9 \rightarrow 5 \rightarrow 0$ จำนวน 48 หน่วย

เส้นทางที่ 2 จาก $0 \rightarrow 2 \rightarrow 3 \rightarrow 6 \rightarrow 0$ จำนวน 39 หน่วย

เส้นทางที่ 3 จาก $0 \rightarrow 1 \rightarrow 7 \rightarrow 0$ จำนวน 35 หน่วย

โดยมีระยะทางรวมเท่ากับ 488 หน่วย

โดยนำแต่ละเส้นทางมาทำการลากเส้นทางได้ผลลัพธ์แสดงดังภาพที่ 8.4



ภาพที่ 8.4 ผลลัพธ์การจัดเส้นทาง

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

8.5 การจัดเส้นทางขนส่งด้วยการเรียนรู้ของเครื่อง

วิธีการเรียนรู้ของเครื่องที่ใช้ในการจัดเส้นทางขนส่ง (VRP) ได้แก่ วิธี การจัดกลุ่มเชิงสเปกตรัม (Spectral Clustering, SC) โดยใช้วิธี SC เพื่อแบ่งลูกค้าออกเป็นกลุ่มที่อยู่ในพื้นที่ใกล้เคียงกัน ช่วยให้สามารถจัดเส้นทางเดินรถได้ง่ายขึ้น และสามารถลดความซับซ้อนของปัญหาได้นอกจากนั้น วิธี SC เป็นเครื่องมือที่มีประสิทธิภาพในการจัดกลุ่มข้อมูลที่ซับซ้อนแม้ว่าข้อมูลจะมี

โครงสร้างที่ไม่ใช่เชิงเส้น และสามารถนำไปใช้ได้หลากหลายรูปแบบข้อมูล เช่น ข้อมูลเชิงเครือข่าย ข้อมูลเชิงกราฟ รายละเอียดโดยสรุปขั้นตอนของวิธี SC อธิบายดังรายละเอียดดังต่อไปนี้

1. สร้างกราฟจากข้อมูล (Graph Construction)

ข้อมูลทั้งหมด ไม่ว่าจะเป็นข้อมูลลูกค้าหรือข้อมูลอื่น ๆ จะถูกจัดเก็บในรูปแบบของโหนด (nodes) ภายในกราฟ เส้นเชื่อม (arcs) ที่เชื่อมต่อระหว่างโหนดแต่ละตัวใน VRP จะแสดงให้เห็นถึงความสัมพันธ์ระหว่างข้อมูลของลูกค้า โดยค่าของเส้นเชื่อมนั้นมักจะแทนด้วยระยะทางระหว่างลูกค้าแต่ละราย ซึ่งจะถูกนำมาใช้ในการคำนวณเส้นทางการส่งสินค้าให้มีประสิทธิภาพสูงสุด ค่าระยะทางระหว่างโหนดถูกนำมาใช้ในการสร้างเมทริกซ์ความใกล้เคียง (affinity matrix) หรือเมทริกซ์ความคล้ายคลึง (similarity matrix) เพื่อแสดงค่าความคล้ายคลึงกันระหว่างโหนดแต่ละคู่ ค่าในเมทริกซ์จะถูกนำไปใช้ในการวิเคราะห์คลัสเตอร์ (clustering analysis) เพื่อจัดกลุ่มโหนดที่มีลักษณะคล้ายคลึงกันเข้าด้วยกัน เมทริกซ์ความใกล้เคียงแสดงดังสมการที่ (8.7)

$$A_{ij} = \exp\left(-\frac{d_{ij}^2}{2\sigma^2}\right) \quad (8.7)$$

โดยที่ A_{ij} แทนค่าความคล้ายคลึงระหว่างโหนด i และ j

d_{ij} แทนระยะทางระหว่างโหนด i และ j

σ แทนพารามิเตอร์ที่ควบคุมความเข้มของความสัมพันธ์

2. คำนวณ Laplacian Matrix

ภายหลังการสร้างเมทริกซ์ความใกล้เคียง (A_{ij}) เสร็จสิ้น ขั้นตอนต่อไปคือการคำนวณเมทริกซ์-ลาปลาเซียน (Laplacian matrix) ซึ่งเปรียบเสมือนแผนที่ที่ช่วยให้เราเข้าใจโครงสร้างและความเชื่อมโยงของข้อมูลในกราฟได้อย่างลึกซึ้งยิ่งขึ้น โดยสามารถคำนวณโดยใช้สมการที่ 8.8 - 8.10

Degree Matrix (D) คือ เมทริกซ์ที่แสดง ดีกรี หรือ จำนวนการเชื่อมต่อ ของแต่ละโหนดในกราฟ โดยค่าในตำแหน่ง D_{ii} จะเท่ากับ ผลรวมของค่าความคล้ายคลึง ทั้งหมดที่เชื่อมต่อกับโหนดที่ i

$$D_{ii} = \sum_j A_{ij} \quad (8.8)$$

Laplacian Matrix (L) คำนวณได้จาก:

$$L = D - A \quad (8.9)$$

หรือในบางกรณีอาจใช้ normalized Laplacian matrix แสดงดังสมการที่ 8.10

$$L_{Normalized} = I - D^{-1/2}AD^{-1/2} \quad (8.10)$$

ซึ่งจะช่วยในกรณีที่กราฟมีโหนดที่มีการเชื่อมต่อกันไม่สม่ำเสมอ

3. การแยกตัวประกอบค่าเอกลักษณะ (Eigenvalue Decomposition)

หลังจากสร้าง Laplacian matrix ขึ้นมาแล้ว ขั้นตอนต่อไปคือการคำนวณหาค่า Eigenvalues และ Eigenvectors เพื่อทำความเข้าใจโครงสร้างกลุ่มภายในข้อมูลที่ซับซ้อน โดย Eigenvectors ที่สอดคล้องกับค่า Eigenvalues ที่มีค่าน้อยที่สุด (ยกเว้นค่าศูนย์) นั้นมีความสำคัญเป็นพิเศษ เพราะจะช่วยให้เรา ระบุกลุ่มย่อยที่เชื่อมโยงกันอย่างใกล้ชิดภายในข้อมูล ได้อย่างแม่นยำ

การวิเคราะห์ Eigenvectors เหล่านี้ไปเปรียบเสมือนการ ลดมิติของข้อมูล ทำให้เราสามารถมองเห็นรูปแบบและความสัมพันธ์ที่ซ่อนอยู่ภายในข้อมูลที่มีมิติสูงได้ชัดเจนยิ่งขึ้น ผลลัพธ์ที่ได้จะช่วยให้เราสามารถ แบ่งข้อมูลออกเป็นกลุ่มที่แตกต่างกัน ได้อย่างง่ายดายในปริภูมิที่มีมิติต่ำกว่า

4. จัดกลุ่มข้อมูลในพื้นที่ใหม่ (Clustering in the Reduced Space)

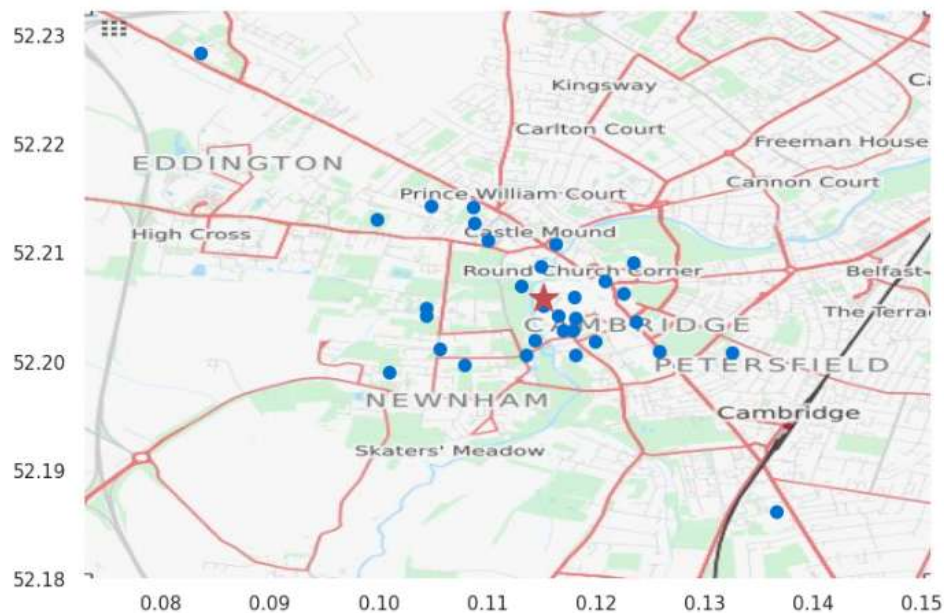
หลังจากที่เราสร้าง 'เมทริกซ์ Laplacian' ซึ่งเป็นเหมือนแผนที่แสดงความเชื่อมโยงระหว่างข้อมูลแต่ละจุดได้แล้ว ขั้นตอนต่อไปคือการหา Eigenvector ซึ่งเป็นเหมือนแกนหลัก (ที่มีค่า Eigenvalue ต่ำที่สุด) ที่สามารถอธิบายข้อมูลทั้งหมดได้ โดยเราจะเลือกแกนหลักที่สำคัญที่สุดมาใช้ในการลดมิติของข้อมูล การลดมิติ หมายถึงการเปลี่ยนข้อมูลที่มีมิติมาก ให้เหลือมิติน้อยลงแต่ยังคงเก็บข้อมูลสำคัญไว้ได้มากที่สุด พุดง่าย ๆ ก็คือ การเปลี่ยนข้อมูลที่ซับซ้อนให้เป็นข้อมูลที่เรียบง่ายขึ้น แต่ยังคงเห็นภาพรวมได้ชัดเจน เมื่อเราได้ข้อมูลในมิติใหม่แล้ว เราจะนำไปใช้กับ อัลกอริทึมการจัดกลุ่ม เช่น K-means ซึ่งเป็นเหมือนเครื่องมือในการแบ่งข้อมูลออกเป็นกลุ่มๆ ที่มีความคล้ายคลึงกัน ตัวอย่างเช่น เราอาจจะต้องการแบ่งลูกค้าออกเป็นกลุ่มตามพฤติกรรมการซื้อสินค้า หรือแบ่งภาพออกเป็นกลุ่มตามลักษณะของวัตถุในภาพ

5. ผลลัพธ์ที่ได้จาก (Spectral Clustering)

SC ให้ผลลัพธ์เป็นกลุ่มข้อมูลที่สอดคล้องกับโครงสร้างที่ซับซ้อนของข้อมูลได้ดีกว่าการจัดกลุ่มแบบดั้งเดิม เช่น K-means ซึ่งมักจะสร้างกลุ่มที่มีขอบเขตเป็นเส้นตรง SC สามารถจับความสัมพันธ์ที่ไม่เป็นเชิงเส้นระหว่างข้อมูลได้ ทำให้สามารถแบ่งกลุ่มข้อมูลที่มีรูปร่างซับซ้อน เช่น กลุ่มข้อมูลที่มีรูปร่างเป็นวงแหวนหรือรูปทรงคล้ายดัมเบลล์ ได้อย่างแม่นยำ

ตัวอย่าง 8.5 การจัดเส้นทางด้วยวิธี Spectral Clustering

เนื่องจากวิธี SC มีความยากในการคำนวณด้วยมือ ดังนั้นในตัวอย่างนี้จะแสดงการประยุกต์ใช้ SC สำหรับปัญหา VRP โดยใช้โปรแกรม Python เข้ามาช่วยในการทำ โดยตัวอย่างได้ดัดแปลงมาจาก (Liu, 2022) ปัญหานี้กำหนดให้มีรถขนส่งจำนวน 4 คัน และจะต้องทำการขนส่งอาหารให้กับมหาวิทยาลัยที่อยู่ในเมืองเคมบริดจ์ทั้งหมด 31 มหาวิทยาลัย ข้อมูลพิกัดและตำแหน่งที่ตั้งของมหาวิทยาลัยแสดงดังภาพที่ 8.5



ภาพที่ 8.5 พิกัดของมหาวิทยาลัยที่ต้องไปส่งอาหาร (หมายเหตุ: รูปดาวคือตำแหน่งร้านอาหาร)

ที่มา: ดัดแปลงจาก: (Liu, 2022)

จากมหาวิทยาลัยเราสามารถเรียกใช้วิธี SC โดยมี Code โปรแกรม Python ดังต่อไปนี้

```
# Import library sklearn
from sklearn.cluster import SpectralClustering

# Input colleges data
colleges = pd.read_excel('Colleges.xlsx', index_col="College" )
colleges
location = Cambridge[['Longitude', 'Latitude']].drop(index='The Pizza Store')

# Initialize the model
v = 4 # number of delivery vehicles
clt = SpectralClustering(n_clusters=v, random_state=42,
                        affinity='nearest_neighbors')
clt.fit_predict(colleges)

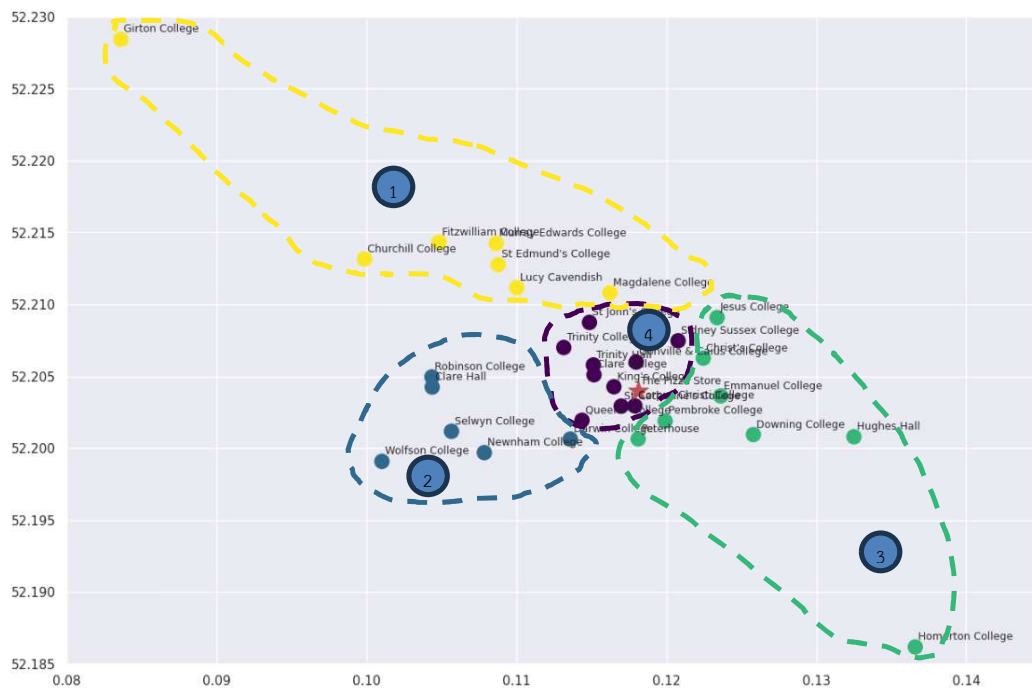
# Plot colleges
fig, ax = plt.subplots(figsize=(15, 10))
ax.scatter(location.Longitude, location.Latitude,
           c=clt.labels_, cmap='viridis',s=150)

# Annotation
for txt in Cambridge.index:
    ax.annotate(txt, xy=(Cambridge.loc[txt].Longitude+0.0002,
                        Cambridge.loc[txt].Latitude+0.0005), size=9)

# Plot the pizza store
ax.scatter(Cambridge.iloc[0].Longitude,
           Cambridge.iloc[0].Latitude, marker="x", c='r', s=300);

# Set x,y limits
ax.set_xlim(0.080, 0.145)
ax.set_ylim(52.185, 52.230);
```

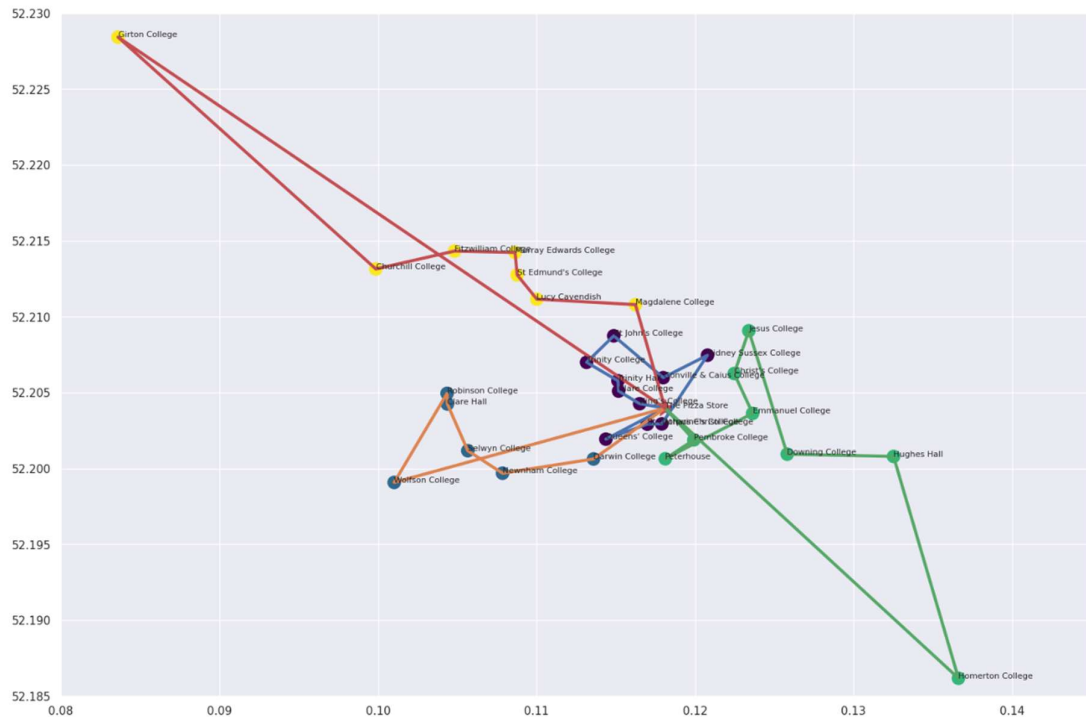
เมื่อทำการจัดกลุ่มโดยใช้วิธี SC จะสามารถแบ่งกลุ่มเส้นทางออกเป็น 4 กลุ่มใหญ่ แสดงดังภาพที่ 8.6 โดยที่เราสามารถทำการจัดเส้นทางขนส่งในแต่ละกลุ่มโดยใช้วิธีการต่อหนดที่ใกล้ที่สุด (nearest neighbor, NN) เส้นทางที่จัดได้ในแต่ละกลุ่มแสดงดังภาพที่ 8.7



ภาพที่ 8.6 กลุ่มเส้นทางหลังจากทำ Spectral Clustering

(หมายเหตุ: รูปดาวคือตำแหน่งร้านอาหาร)

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)



ภาพที่ 8.7 เส้นทางขนส่ง (หมายเหตุ: รูปดาวคือตำแหน่งร้านอาหาร)

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

จากภาพที่ 8.7 จะเห็นได้ว่าเราสามารถจัดเส้นทางได้ 4 เส้นทางตามยานพาหนะที่มีอยู่ 4 คัน เส้นทางโดยสรุปแสดงดังต่อไปนี้

เส้นทางที่ 1 สีเหลือง The Pizza Store → Girton College → Churchill College → Fitzwilliam College → Murray Edwards College → St Edmund's College → Lucy Cavendish → Magdalene College → The Pizza Store

เส้นทางที่ 2 สีฟ้า The Pizza Store → Wolfson College → Clare Hall → Robinson College → Selwyn College → Newnham College → Darwin College → The Pizza Store

เส้นทางที่ 3 สีม่วง The Pizza Store → Trinity College → St John's College → Trinity Hall → Clare College → King's College → Gonville & Caius College → Corpus Christi College → Queens' College → Pembroke College → The Pizza Store

เส้นทางที่ 4 สีเขียว The Pizza Store → Jesus College → Sidney Sussex College → Christ's College → Emmanuel College → Peterhouse → Downing College → Hughes Hall → Homerton College → The Pizza Store

สรุป

การจัดเส้นทางขนส่ง (VRP) เป็นหัวใจสำคัญของธุรกิจโลจิสติกส์ที่ต้องการความได้เปรียบในการแข่งขัน เพราะการจัดเส้นทางที่มีประสิทธิภาพจะช่วยให้ธุรกิจมีข้อได้เปรียบหลายประการ ได้แก่ (1) ลดค่าใช้จ่ายด้านเชื้อเพลิง บำรุงรักษา และค่าแรงได้อย่างมีนัยสำคัญ (2) เพิ่มความรวดเร็วและความแม่นยำในการส่งมอบ เพราะสินค้าถึงมือลูกค้าตรงเวลาและครบถ้วน (3) เพิ่มประสิทธิภาพการใช้ทรัพยากร เนื่องจากใช้รถขนส่งน้อยลง ลดการปล่อยมลพิษ และลดความเสียหายต่อสิ่งแวดล้อม (4) ปรับตัวเข้ากับสถานะที่เปลี่ยนแปลงได้อย่างรวดเร็ว เมื่อเกิดเหตุการณ์ไม่คาดคิด เช่น อุบัติเหตุหรือการจราจรติดขัด ระบบจัดเส้นทางที่ทันสมัยสามารถปรับเปลี่ยนเส้นทางได้ทันที ทำให้การส่งมอบสินค้าไม่ล่าช้า และ (6) เพิ่มขีดความสามารถในการแข่งขัน เพราะธุรกิจที่มีระบบการจัดเส้นทางที่ดี จะสามารถให้บริการลูกค้าได้อย่างรวดเร็ว แม่นยำ และตรงตามความต้องการ ซึ่งเป็นปัจจัยสำคัญในการสร้างความแตกต่างและเอาชนะคู่แข่งได้ อย่างไรก็ตาม ปัญหาการจัดเส้นทางขนส่งเป็นปัญหาที่ยากซับซ้อนและใช้เวลาในการหาคำตอบด้วยคอมพิวเตอร์เป็นระยะเวลานานเมื่อลูกค้ามีจำนวนมาก โดยเฉพาะเมื่อต้องการคำตอบที่เหมาะสมที่สุด (optimal solutions) เช่น ถ้าลูกค้ามีจำนวน 20 ราย จะมีเส้นทางที่เป็นไปได้ทั้งหมด $1.21645100408832 \times 10^{18}$ เส้นทาง ทำให้ถึงแม้ใช้คอมพิวเตอร์ช่วยอาจใช้เวลาหลายวันในการคำนวณ ดังนั้นในการแก้ปัญหาด้วยคอมพิวเตอร์ เมื่อต้องการคำตอบที่ดีที่สุดนั้นจะสามารถทำได้เฉพาะปัญหาที่มีขนาดเล็กโดยใช้เครื่องมืออย่างเช่น โปรแกรมเชิงเส้น (LP) แต่สำหรับปัญหาขนาดใหญ่เพื่อลดเวลาในการคำนวณ นักวิจัยหลายคนได้พัฒนาวิธีการหาคำตอบที่ใกล้เคียงคำตอบที่ดีที่สุด (near optimal solutions) ด้วยเวลาในการคำนวณที่สั้นและรวดเร็ว ซึ่งวิธีเหล่านี้ถูกเรียกรวมๆ กันว่า วิธีฮิวริสติก (heuristic) โดยในบทนี้ได้กล่าวถึงวิธีฮิวริสติกวิธีหนึ่งที่เป็นที่นิยมใช้ได้แก่ วิธี saving algorithm ของ (Clarke & Wright, 1964) เพราะเป็นวิธีที่ง่ายต่อการใช้งานและสามารถลดเวลาในการคำนวณหาคำตอบได้อย่างมีนัยสำคัญ ด้วยเหตุนี้เองจึงได้มีการนำวิธี Saving มาเพิ่มพารามิเตอร์อย่างเช่น ช่วงเวลาในการขนส่งได้ (time windows) รถบรรทุกที่มีขนาดและความจุแตกต่างกัน (Heterogeneous fleet) หรือ อื่นๆ โดยได้มีการเขียนโปรแกรมคอมพิวเตอร์ประกอบการใช้งานที่สะดวกรวดเร็วและสามารถเข้าถึงได้บนคลาวด์

แบบฝึกหัด

1. บริษัท XYZ ต้องการจัดส่งสินค้าจากคลังสินค้า A ไปยังจุดหมายปลายทาง B, C, และ D โดยมีเส้นทางเชื่อมต่อระหว่างแต่ละจุดหลายเส้นทาง คุณได้รับข้อมูลเส้นทางดังต่อไปนี้:

A ↔ B: ระยะทาง 10 กิโลเมตร

A ↔ C: ระยะทาง 15 กิโลเมตร

A ↔ D: ระยะทาง 20 กิโลเมตร

B ↔ C: ระยะทาง 12 กิโลเมตร

B ↔ D: ระยะทาง 8 กิโลเมตร

C ↔ D: ระยะทาง 7 กิโลเมตร

จงวิเคราะห์และคำนวณเส้นทางที่เป็นไปได้ทั้งหมดที่สามารถจัดส่งสินค้าได้ โดยพิจารณาทั้งเส้นทางไปกลับและไม่ย้อนเส้นทางเดิม และให้ระบุเส้นทางที่สั้นที่สุดในการขนส่งสินค้าครบทุกจุด

2. จงอธิบายถึงปัญหาการจัดเส้นทางขนส่ง (Vehicle Routing Problem, VRP) ว่ามีลักษณะอย่างไร และจงยกตัวอย่างสถานการณ์ที่ปัญหานี้สามารถนำไปประยุกต์ใช้ในธุรกิจได้
3. ในกรณีที่มีศูนย์กระจายสินค้าอยู่หลายแห่ง และลูกค้ามีความต้องการสินค้าที่แตกต่างกันไปตามสถานที่ต่างๆ จงอธิบายวิธีการใช้แบบจำลองทางคณิตศาสตร์ในการหาวิธีการจัดเส้นทางขนส่งที่มีต้นทุนต่ำที่สุด
4. ปัญหา Traveling Salesman Problem (TSP) เป็นปัญหาพื้นฐานที่พบในงานจัดเส้นทางขนส่ง จงอธิบายถึงปัญหานี้และวิธีการแก้ไขปัญหานี้ด้วยวิธีการเชิงฮิวริสติก (Heuristic Approach)
5. จงเปรียบเทียบข้อดีและข้อเสียของการใช้ Genetic Algorithm (GA) และ Particle Swarm Optimization (PSO) ในการหาวิธีการจัดเส้นทางขนส่งที่เหมาะสมที่สุด
6. การใช้เส้นทางหลายเส้นทาง (Multiple Routes) ในการจัดส่งสินค้าไปยังปลายทางต่างๆ จะสามารถลดต้นทุนการขนส่งได้อย่างไร จงอธิบายพร้อมยกตัวอย่างสถานการณ์จริงที่มีการนำแนวคิดนี้ไปใช้
7. ในกรณีที่เส้นทางขนส่งมีข้อจำกัดในการเข้าถึง เช่น ถนนมีความแคบหรือมีข้อจำกัดในน้ำหนักบรรทุก จงเสนอวิธีการจัดเส้นทางขนส่งที่เหมาะสมที่สุดพร้อมอธิบายหลักการในการตัดสินใจเลือกเส้นทาง

8. ลูกค้านี้ตำแหน่งที่ตั้ง โหนด 1-9 และ Depot อยู่ที่ตำแหน่ง 0 โดยมีพิกัดและความต้องการของแต่ละโหนดแสดงดังตารางต่อไปนี้ (รถบรรทุกมีความจุ = 50) จงใช้วิธี Saving Algorithm จัดเส้นทางการขนส่งที่เหมาะสม

Node	x	y	Demand
0	42	68	0
1	77	97	5
2	28	64	23
3	77	39	14
4	32	33	13
5	32	8	8
6	42	92	18
7	8	3	19
8	7	14	10
9	82	17	18

เอกสารอ้างอิง

- Applegate, D. L., Bixby, R. E., Chvátal, V., & Cook, W. J. (2006). *The traveling salesman problem: A computational study*. Princeton University Press.
- Bai, R., Chen, X., Chen, Z.-L., Cui, T., Gong, S., He, W., & Zhang, H. (2021). Analytics and machine learning in vehicle routing research: Methods and applications. *International Journal of Production Research*, 59(3), 820-846. <https://doi.org/10.1080/00207543.2021.2013566>.
- Bräysy, O., & Gendreau, M. (2005). Vehicle routing problem with time windows, Part I: Route construction and local search algorithms. *Transportation Science*, 39(1), 104–118. <https://doi.org/10.1287/trsc.1030.0056>
- Braekers, K., Ramaekers, K., & Van Nieuwenhuysse, I. (2016). The vehicle routing problem: State of the art classification and review. *Computers & Industrial Engineering*, 99, 300-313. <https://doi.org/10.1016/j.cie.2015.12.007>.
- Clarke, G., & Wright, J. W. (1964). Scheduling of vehicles from a central depot to a number of delivery points. *Operations Research*, 12(4), 568-581. <https://doi.org/10.1287/opre.12.4.568>.
- Dantzig, G. B., Fulkerson, D. R., & Johnson, S. M. (1954). Solution of a large-scale traveling-salesman problem. *Journal of the Operations Research Society of America*, 2(4), 393–410. <https://doi.org/10.1287/opre.2.4.393>
- Dantzig, G. B., & Ramser, J. H. (1959). The truck dispatching problem. *Management Science*, 6(1), 80–91. <https://doi.org/10.1287/mnsc.6.1.80>
- Fisher, M. L., & Jaikumar, R. (1981). A generalized assignment heuristic for vehicle routing. *Networks*, 11(2), 109-124. <https://doi.org/10.1002/net.3230110205>.
- Gendreau, M., Laporte, G., & Potvin, J.-Y. (2010). Metaheuristics for the capacitated VRP. In P. Toth & D. Vigo (Eds.), *Vehicle routing: Problems, methods, and applications* (2nd ed., pp. 129–154). SIAM.

เอกสารอ้างอิง (ต่อ)

- Gillett, B. E., & Miller, L. R. (1974). A heuristic algorithm for the vehicle-dispatch problem. *Operations Research*, 22(2), 340-349.
<https://doi.org/10.1287/opre.22.2.340>.
- Hillier, F. S., & Lieberman, G. J. (2021). *Introduction to operations research* (11th ed.). McGraw-Hill Education.
- Laporte, G. (2009). Fifty years of vehicle routing. *Transportation Science*, 43(4), 408–416. <https://doi.org/10.1287/trsc.1090.0301>
- Laporte, G., & Semet, F. (2002). Classical heuristics for the capacitated VRP. In P. Toth & D. Vigo (Eds.), *The vehicle routing problem* (pp. 109-128). Society for Industrial and Applied Mathematics (SIAM). <https://doi.org/10.1137/1.9780898718515.ch5>.
- Lawler, E. L., Lenstra, J. K., Rinnooy Kan, A. H. G., & Shmoys, D. B. (1985). *The traveling salesman problem: A guided tour of combinatorial optimization*. Wiley.
- Lenstra, J. K., & Rinnooy Kan, A. H. G. (1981). Complexity of vehicle routing and scheduling problems. *Networks*, 11(2), 221–227.
<https://doi.org/10.1002/net.3230110211>
- Liu, K. Y. (2022). *Supply Chain Analytics: Concepts, Techniques and Applications*. (1st ed.) Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-92224-5>.
- Miller, C. E., Tucker, A. W., & Zemlin, R. A. (1960). Integer programming formulation of traveling salesman problems. *Journal of the ACM*, 7(4), 326-329.
<https://doi.org/10.1145/321043.321046>.
- Osaba, E., Yang, X.-S., Diaz, F., Lopez-Garcia, P., & Carballedo, R. (2020). An improved discrete bat algorithm for symmetric and asymmetric traveling salesman problems. *Engineering Applications of Artificial Intelligence*, 48, 59-71.
<https://doi.org/10.1016/j.engappai.2015.10.006>.

เอกสารอ้างอิง (ต่อ)

Solomon, M. M. (1987). Algorithms for the vehicle routing and scheduling problems with time window constraints. *Operations Research*, 35(2), 254-265.
<https://doi.org/10.1287/opre.35.2.254>.

Toth, P., & Vigo, D. (Eds.). (2014). *Vehicle routing: Problems, methods, and applications* (2nd ed.). SIAM.



ส่วนที่ 5

กรณีศึกษา/งานวิจัยที่เกี่ยวข้องและน่าสนใจ

Part 5: Related and Interested
Case Studies/Research

บทที่ 9

การประยุกต์ใช้การเรียนรู้ของเครื่องในการจัดกลุ่มลูกค้า Applying Machine Learning to Customer Segmentation

หัวข้อ

- 9.1 วัตถุประสงค์ของกรณีศึกษา
- 9.2 บริบทธุรกิจ
- 9.3 ชุดข้อมูลและตัวแปรที่ใช้
- 9.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง
- 9.5 สรุปผลลัพธ์จากการจัดกลุ่มลูกค้า

ในยุคที่ข้อมูลมีปริมาณเพิ่มขึ้นอย่างรวดเร็วและมีความหลากหลายสูง องค์กรธุรกิจ โดยเฉพาะในภาคการค้าปลีกและพาณิชย์อิเล็กทรอนิกส์ จำเป็นต้องพัฒนาความสามารถในการวิเคราะห์ข้อมูลเพื่อทำความเข้าใจพฤติกรรมลูกค้าอย่างลึกซึ้ง การจัดกลุ่มลูกค้า (customer segmentation) จึงเป็นเครื่องมือสำคัญที่ช่วยให้องค์กรสามารถจำแนกลูกค้าออกเป็นกลุ่มที่มีลักษณะหรือพฤติกรรมคล้ายคลึงกัน และนำไปสู่การออกแบบกลยุทธ์ที่มีประสิทธิภาพมากขึ้น (Ngai et al., 2009)

ในอดีต การจัดกลุ่มลูกค้ามักอาศัยวิธีการเชิงสถิติพื้นฐานหรือการแบ่งกลุ่มตามลักษณะประชากรศาสตร์ เช่น อายุ เพศ หรือรายได้ อย่างไรก็ตาม วิธีการดังกล่าวมีข้อจำกัดในการสะท้อนพฤติกรรมที่แท้จริงของลูกค้า งานวิจัยด้านการทำเหมืองข้อมูล (data mining) และการเรียนรู้ของเครื่อง (ML) ได้ชี้ให้เห็นว่า เทคนิคการจัดกลุ่มข้อมูล (clustering) สามารถค้นหารูปแบบที่ซ่อนอยู่ในข้อมูลได้อย่างมีประสิทธิภาพ โดยเฉพาะอัลกอริทึม เช่น K-means ซึ่งเป็นหนึ่งในวิธีที่ได้รับความนิยมสูงและถูกใช้อย่างแพร่หลาย (Jain, 2010; Chen et al., 2012)

ในบริบทของการจัดการซัพพลายเชน การจัดกลุ่มลูกค้าไม่ได้จำกัดอยู่เพียงด้านการตลาดเท่านั้น แต่ยังมีบทบาทสำคัญในการสนับสนุนการตัดสินใจเชิงกลยุทธ์และเชิงปฏิบัติการ เช่น การพยากรณ์อุปสงค์ การบริหารสินค้าคงคลัง และการออกแบบเครือข่ายการกระจายสินค้า โดยการใช้ข้อมูลลูกค้าร่วมกับเทคนิคการวิเคราะห์ขั้นสูงสามารถช่วยเพิ่มความแม่นยำและความยืดหยุ่นของซัพพลายเชนได้อย่างมีนัยสำคัญ (Waller & Fawcett, 2013) ในช่วงหลัง การพัฒนาเทคนิค

ปัญญาประดิษฐ์และการเรียนรู้เชิงลึก (deep learning) ได้เปิดโอกาสใหม่ในการจัดกลุ่มลูกค้าจากข้อมูลขนาดใหญ่และซับซ้อน โดยวิธีการ เช่น deep embedded clustering สามารถเรียนรู้โครงสร้างของข้อมูลและทำการจัดกลุ่มได้พร้อมกัน ส่งผลให้มีประสิทธิภาพสูงกว่าวิธีแบบดั้งเดิมในหลายกรณี (Guo et al., 2017)

บทนี้มีวัตถุประสงค์เพื่ออธิบายแนวคิดและวิธีการจัดกลุ่มลูกค้าด้วยเทคนิคการเรียนรู้ของเครื่อง โดยเชื่อมโยงทั้งในเชิงทฤษฎีและการประยุกต์ใช้จริง ผ่านกรณีศึกษาที่ใช้ภาษา Python ในการวิเคราะห์ข้อมูลลูกค้า ผู้อ่านจะได้เข้าใจทั้งหลักการพื้นฐาน ขั้นตอนการดำเนินงาน และแนวทางการนำผลลัพธ์ไปใช้ในการตัดสินใจด้านซัพพลายเชนอย่างเป็นระบบ ซึ่งเป็นพื้นฐานสำคัญสำหรับการพัฒนาโซลูชันด้านการวิเคราะห์ข้อมูลลูกค้า (customer analytics) ในระดับที่สูงขึ้น

9.1 วัตถุประสงค์ของกรณีศึกษา

กรณีศึกษานี้มีวัตถุประสงค์เพื่อแสดงการประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่องในการจัดกลุ่มลูกค้าในบริบทธุรกิจค้าปลีกออนไลน์ โดยมุ่งเน้นการพัฒนากรอบการวิเคราะห์ที่สามารถสะท้อนพฤติกรรมลูกค้าและความแตกต่างเชิงมูลค่าทางธุรกิจได้อย่างมีประสิทธิภาพ ทั้งนี้ งานวิจัยต้นแบบได้ชี้ให้เห็นว่าการแบ่งกลุ่มลูกค้าควรพิจารณาทั้งมิติพฤติกรรมและมิติทางเศรษฐศาสตร์ร่วมกัน เพื่อให้ผลลัพธ์สามารถนำไปใช้เชิงกลยุทธ์ได้จริง

โดยวัตถุประสงค์ของกรณีศึกษานี้คือ (1) เพื่อพัฒนาแบบจำลองการจัดกลุ่มลูกค้าด้วยเทคนิค ML โดยใช้ข้อมูลพฤติกรรมการซื้อ เช่น Recency Frequency และ Monetary (RFM) (2) เพื่อเปรียบเทียบผลการจัดกลุ่มระหว่างวิธีดั้งเดิม (RFM scoring) และวิธีการเรียนรู้ของเครื่อง (K-Means clustering) รวมไปถึงการทำความเข้าใจความแตกต่างของโครงสร้างกลุ่มลูกค้า (3) เพื่อวิเคราะห์ลักษณะและมูลค่าของลูกค้าในแต่ละกลุ่ม เพื่อสนับสนุนการกำหนดกลยุทธ์ทางธุรกิจและการบริหารลูกค้าอย่างมีประสิทธิภาพ

9.2 บริบทธุรกิจ

ธุรกิจค้าปลีกออนไลน์ในปัจจุบันมีการเติบโตอย่างรวดเร็วจากการขยายตัวของพาณิชย์อิเล็กทรอนิกส์ (E-commerce) ส่งผลให้องค์กรสามารถเข้าถึงข้อมูลธุรกรรมของลูกค้าในระดับที่ละเอียดและมีปริมาณมาก ข้อมูลดังกล่าวครอบคลุมพฤติกรรมการซื้อสินค้า เช่น ความถี่ในการซื้อ มูลค่าการใช้จ่าย และช่วงเวลาการทำธุรกรรม ซึ่งสะท้อนลักษณะพฤติกรรมของลูกค้าได้อย่างชัดเจน ในกรณีศึกษานี้ องค์กรเป็นธุรกิจค้าปลีกออนไลน์ที่มีข้อมูลธุรกรรมจำนวนมากจากลูกค้าหลายพื้นที่ โดยข้อมูลประกอบด้วยรายละเอียด เช่น รายการสั่งซื้อ จำนวนสินค้า ราคาต่อหน่วย วันที่ทำรายการ

และรหัสลูกค้า ซึ่งสามารถนำมาประมวลผลเพื่อสร้างตัวแปรเชิงพฤติกรรม เช่น Recency Frequency และ Monetary (RFM) สำหรับใช้ในการวิเคราะห์ลูกค้า

อย่างไรก็ตาม ความท้าทายสำคัญของธุรกิจคือ “ความหลากหลายของพฤติกรรมลูกค้า” (customer heterogeneity) ซึ่งทำให้ไม่สามารถใช้กลยุทธ์เดียวตอบสนองลูกค้าทุกกลุ่มได้อย่างมีประสิทธิภาพ ลูกค้าบางกลุ่มมีความถี่ในการซื้อสูงและสร้างรายได้หลักให้กับองค์กร ขณะที่บางกลุ่มอาจมีแนวโน้มลดการซื้อหรือเลิกใช้บริการ ซึ่งส่งผลโดยตรงต่อการพยากรณ์อุปสงค์และการบริหารสินค้าคงคลัง ในอดีต องค์กรมักใช้วิธีการแบ่งกลุ่มลูกค้าแบบดั้งเดิม เช่น การวิเคราะห์ RFM ซึ่งมีข้อดีในด้านความเข้าใจง่ายและสามารถนำไปใช้ได้ทันที อย่างไรก็ตาม วิธีดังกล่าวมีข้อจำกัดในการสะท้อนความซับซ้อนของพฤติกรรมลูกค้าในยุคดิจิทัล โดยเฉพาะเมื่อข้อมูลมีหลายมิติและมีขนาดใหญ่ ดังนั้น องค์กรจึงมีความจำเป็นต้องนำเทคนิค ML มาใช้ในการจัดกลุ่มลูกค้า เพื่อค้นหารูปแบบที่ซ่อนอยู่ในข้อมูลและสร้างความเข้าใจเชิงลึกเกี่ยวกับลูกค้าแต่ละกลุ่ม การใช้วิธีการดังกล่าวไม่เพียงช่วยเพิ่มความแม่นยำในการแบ่งกลุ่มลูกค้า แต่ยังช่วยสนับสนุนการตัดสินใจเชิงกลยุทธ์ในด้านต่าง ๆ เช่น การออกแบบแคมเปญการตลาด การรักษาลูกค้า และการเพิ่มประสิทธิภาพของซัพพลายเชน

โดยเฉพาะในบริบทของการจัดการซัพพลายเชน การเข้าใจพฤติกรรมของลูกค้าแต่ละกลุ่มมีความสำคัญต่อการพยากรณ์อุปสงค์ การวางแผนสินค้าคงคลัง และการกำหนดระดับการให้บริการที่เหมาะสม ซึ่งจะช่วยให้องค์กรสามารถลดต้นทุนและเพิ่มความสามารถในการแข่งขันในระยะยาว

9.3 ชุดข้อมูลและตัวแปรที่ใช้

กรณีศึกษาที่ใช้ข้อมูลธุรกรรมของลูกค้า (transactional data) จากธุรกิจค้าปลีกออนไลน์ ซึ่งเป็นข้อมูลระดับรายการสั่งซื้อ (transaction-level data) ที่สะท้อนพฤติกรรมการซื้อของลูกค้าในช่วงเวลาที่กำหนด ข้อมูลดังกล่าวมีปริมาณมากและครอบคลุมลูกค้าจากหลายพื้นที่ ทำให้เหมาะสำหรับการวิเคราะห์เชิงพฤติกรรมและการประยุกต์ใช้เทคนิค ML

9.3.1 รายละเอียดของชุดข้อมูล

ชุดข้อมูลที่ใช้ในกรณีศึกษานี้เป็นข้อมูลธุรกรรมของลูกค้าในธุรกิจค้าปลีกออนไลน์ โดยครอบคลุมช่วงเวลาตั้งแต่เดือนเมษายน พ.ศ. 2567 ถึงเดือนเมษายน พ.ศ. 2568 ซึ่งเป็นข้อมูลล่าสุดที่สะท้อนพฤติกรรมลูกค้าในบริบทของพาณิชย์อิเล็กทรอนิกส์ในปัจจุบัน

9.3.2 การแปลงข้อมูลเชิงพฤติกรรม (Feature Construction)

เพื่อให้สามารถนำข้อมูลไปใช้ในการจัดกลุ่มลูกค้าได้อย่างมีประสิทธิภาพ ข้อมูลธุรกรรมจะถูกแปลงให้อยู่ในระดับลูกค้า (customer-level aggregation) และสร้างตัวแปรเชิงพฤติกรรมที่สำคัญ

ได้แก่ Recency (R) ระยะเวลาตั้งแต่การซื้อครั้งล่าสุดของลูกค้า Frequency (F) จำนวนครั้งที่ลูกค้าทำธุรกรรม Monetary (M) มูลค่าการใช้จ่ายรวมของลูกค้า ตัวแปร RFM เป็นตัวแปรพื้นฐานที่ได้รับ การยอมรับอย่างแพร่หลายในการวิเคราะห์ลูกค้า เนื่องจากสามารถสะท้อนทั้ง “ความใหม่ของการซื้อ” “ความถี่ในการซื้อ” และ “มูลค่าทางเศรษฐกิจ” ของลูกค้าได้อย่างครอบคลุม

9.3.3 การเตรียมข้อมูลสำหรับ Machine Learning

เนื่องจากตัวแปร RFM และตัวแปรอื่น ๆ มีหน่วยวัดที่แตกต่างกัน เช่น จำนวนวัน จำนวนครั้ง และจำนวนเงิน จึงจำเป็นต้องมีการปรับสเกลข้อมูล (data scaling) เพื่อให้สามารถนำไปวิเคราะห์ ร่วมกันได้อย่างเหมาะสม โดยอาจใช้วิธี เช่น Standardization หรือ Log Transformation (ในกรณี ข้อมูลกระจายไม่สมมาตร) ขั้นตอนดังกล่าวช่วยลดอิทธิพลของตัวแปรที่มีค่าสูงมาก และทำให้โมเดล clustering สามารถเรียนรู้โครงสร้างของข้อมูลได้ดีขึ้น

ชุดข้อมูลที่ใช้ในกรณีศึกษานี้เป็นข้อมูลธุรกรรมจริงที่ถูกแปลงให้อยู่ในรูปแบบเชิงพฤติกรรม ของลูกค้า ซึ่งเหมาะสมสำหรับการประยุกต์ใช้เทคนิค ML โดยเฉพาะการจัดกลุ่มลูกค้า เนื่องจาก สามารถสะท้อนความแตกต่างของลูกค้าในหลายมิติ และรองรับการวิเคราะห์เชิงลึกในระดับธุรกิจ และซัพพลายเชน

9.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง

สำหรับหัวข้อนี้จะแสดงขั้นตอนการวิเคราะห์แบ่งกลุ่มลูกค้าโดยละเอียด ซึ่งจะมีการ สอดแทรกการเขียนโค้ดโดยใช้โปรแกรม PYTHON ซึ่งเหมาะสมกับปริมาณของข้อมูลที่มีประมาณ 541,390 รายการ เพื่อไม่ให้เนื้อหายาวเกินไปจึงจะทำการแสดงบางส่วนที่จำเป็นในหัวข้อนี้

เริ่มจากการนำเข้าข้อมูลและทำการทำความสะอาดข้อมูลโดยใช้ไฟล์ “inputdata.xlsx”

ตัวอย่าง 9.1 ป้อนข้อมูลนำเข้า

```
import pandas as pd

# Load the 'inputdata.xlsx' file into a pandas DataFrame
df_input = pd.read_excel('inputdata.xlsx')
print("File 'inputdata.xlsx' loaded successfully. Here are the first 5 rows:")
display(df_input.head())
```

บทที่ 9 การประยุกต์ใช้การเรียนรู้ของเครื่องในการจัดกลุ่มลูกค้า

File 'inputdata.xlsx' loaded successfully. Here are the first 5 rows:

	InvoiceNo	StockCode	Description	Quantity	ThaiDate	Thai unit price	CustomerID	Province
0	578459	22338	STAR DECORATION PAINTED ZINC	96	2025-04-11 12:30:00	6.27	12388.0	Suphan Buri
1	578459	22600	CHRISTMAS RETROSPOT STAR WOOD	20	2025-04-11 12:30:00	28.05	12388.0	Suphan Buri
2	578459	22910	PAPER CHAIN KIT VINTAGE CHRISTMAS	20	2025-04-11 12:30:00	97.35	12388.0	Suphan Buri
3	578459	22086	PAPER CHAIN KIT 50'S CHRISTMAS	20	2025-04-11 12:30:00	97.35	12388.0	Suphan Buri
4	578459	22340	NOEL GARLAND PAINTED ZINC	24	2025-04-11 12:30:00	12.87	12388.0	Suphan Buri

จากที่แสดงข้างบนจะเห็นได้ว่าข้อมูลประกอบด้วยตัวแปรหลักที่เกี่ยวข้องกับได้แก่

- InvoiceNo: หมายเลขใบสั่งซื้อ ใช้ระบุรายการธุรกรรม
- StockCode: รหัสสินค้า
- Description: รายละเอียดสินค้า
- Quantity: จำนวนสินค้าที่ซื้อในแต่ละรายการ
- InvoiceDate: วันที่และเวลาที่ทำรายการ
- UnitPrice: ราคาต่อหน่วยของสินค้า
- CustomerID: รหัสลูกค้า
- Province: จังหวัดของลูกค้า

ตัวอย่าง 9.2 ทำความสะอาดข้อมูลที่สูญหาย (Missing Value)

```
print(f"Shape of df_input before dropping missing values in CustomerID: {df_input.shape}")
df_input.dropna(subset=["CustomerID"], inplace=True)
print(f"Shape of df_input after dropping missing values in CustomerID: {df_input.shape}")
```

```
Shape of df_input before dropping missing values in CustomerID: (541390, 8)
Shape of df_input after dropping missing values in CustomerID: (406573, 8)
```

จำนวนข้อมูลก่อนทำความสะอาดมีอยู่ 541,390 รายการ หลังจากนั้นข้อมูลดังกล่าวถูกนำมาทำความสะอาด (Data Cleaning) โดยลบรายการที่มีข้อมูลสูญหาย (Missing Values) หรือค่าที่ผิดปกติ เช่น รายการที่ไม่มีรหัสลูกค้า หรือรายการคืนสินค้า เพื่อให้ได้ชุดข้อมูลที่มีความถูกต้องและเหมาะสมสำหรับการวิเคราะห์ ทำให้จำนวนข้อมูลหลังทำความสะอาดเหลืออยู่ 406,573 รายการ ถึงแม้ว่าจะต้องเสียข้อมูลไปบางส่วน แต่จำเป็นที่จะต้องแลกกับผลการวิเคราะห์ที่ถูกต้องแม่นยำ

ตัวอย่าง 9.3 เพิ่มคอลัมน์ TotalPrice เพื่อใช้ในการคำนวณ RFM

```
df_input['TotalPrice'] = df_input['Quantity'] * df_input['Thai unit price']
print("Added 'TotalPrice' column. Here are the first 5 rows with the new column:")
display(df_input.head())
```

Added 'TotalPrice' column. Here are the first 5 rows with the new column:

	InvoiceNo	StockCode	Description	Quantity	ThaiDate	Thai unit price	CustomerID	Province	TotalPrice
0	578459	22338	STAR DECORATION PAINTED ZINC	96	2025-04-11 12:30:00	6.27	12388.0	Suphan Buri	601.92
1	578459	22600	CHRISTMAS RETROSPOT STAR WOOD	20	2025-04-11 12:30:00	28.05	12388.0	Suphan Buri	561.00
2	578459	22910	PAPER CHAIN KIT VINTAGE CHRISTMAS	20	2025-04-11 12:30:00	97.35	12388.0	Suphan Buri	1947.00
3	578459	22086	PAPER CHAIN KIT 50'S CHRISTMAS	20	2025-04-11 12:30:00	97.35	12388.0	Suphan Buri	1947.00
4	578459	22340	NOEL GARLAND PAINTED ZINC	24	2025-04-11 12:30:00	12.87	12388.0	Suphan Buri	308.88

คอลัมน์ TotalPrice ถูกเพิ่มเป็น คอลัมน์สุดท้ายโดยได้มาจากการนำ ปริมาณ (Quantity) x ราคาต่อหน่วย (Unit Price) โดยรูปข้างบนเป็นการแสดงข้อมูลตัวอย่าง 5 แถวแรก หลังจากเพิ่ม TotalPrice แล้ว

ตัวอย่าง 9.4 ทำการคำนวณค่า RFM

การคำนวณ RFM เริ่มจากการคำนวณ R (Recency) (หน่วยเป็นวัน) ของลูกค้าแต่ละราย เป็นตัวแปรแรก โดยดูจากวันที่ซื้อล่าสุดเทียบกับวันที่ปัจจุบันที่ทำการวิเคราะห์ข้อมูล ดังแสดงดังต่อไปนี้

```
analysis_date = df_input['ThaiDate'].max()
print(f"Analysis date: {analysis_date.strftime('%Y-%m-%d %H:%M:%S')}")
```

Analysis date: 2025-04-26 12:50:00

เริ่มจากการหาวันที่ปัจจุบันที่ทำการวิเคราะห์ข้อมูลแสดงดังข้างบน หลังจากนั้นนำ Analysis Date มาลบกับวันที่ซื้อล่าสุดของลูกค้าแต่ละรายแล้วทำการแปลงเป็นจำนวนวันแสดงดังต่อไปนี้

```
last_purchase = df_input.groupby('CustomerID')['ThaiDate'].max().reset_index()
last_purchase.columns = ['CustomerID', 'LastPurchaseDate']

rfm_df = last_purchase
rfm_df['Recency'] = (analysis_date - rfm_df['LastPurchaseDate']).dt.days
```

```
print("Recency (R) calculated successfully. Here are the first 5 rows of the RFM DataFrame:")
display(rfm_df.head())
```

Recency (R) calculated successfully. Here are the first 5 rows of the RFM DataFrame:

	CustomerID	LastPurchaseDate	Recency
0	12346.0	2024-06-05 10:17:00	325
1	12347.0	2025-04-24 15:52:00	1
2	12348.0	2025-02-10 13:13:00	74
3	12349.0	2025-04-08 09:51:00	18
4	12350.0	2024-06-20 16:01:00	309

หลังจากคำนวณ R จะทำการคำนวณตัวแปรถัดไปได้แก่ Frequency (F) โดยเริ่มจากการจัดกลุ่ม CustomerID และนับจำนวน InvoiceNo ที่ไม่เหมือนกันของลูกค้าแต่ละราย แสดงดังต่อไปนี้

```
frequency = df_input.groupby('CustomerID')['InvoiceNo'].nunique().reset_index()
frequency.columns = ['CustomerID', 'Frequency']

rfm_df = rfm_df.merge(frequency, on='CustomerID', how='left')

print("Frequency (F) calculated and merged successfully. Here are the first 5 rows of the RFM DataFrame:")
display(rfm_df.head())
```

Frequency (F) calculated and merged successfully. Here are the first 5 rows of the RFM DataFrame:

	CustomerID	LastPurchaseDate	Recency	Frequency
0	12346.0	2024-06-05 10:17:00	325	2
1	12347.0	2025-04-24 15:52:00	1	7
2	12348.0	2025-02-10 13:13:00	74	4
3	12349.0	2025-04-08 09:51:00	18	1
4	12350.0	2024-06-20 16:01:00	309	1

คอลัมน์ความถี่แสดงดังคอลัมน์สุดท้าย (Frequency) (หน่วยเป็นจำนวน) โดยข้างบนแสดงเฉพาะตัวอย่าง 5 แถวแรก

มาถึงตัวแปรสุดท้ายคือการคำนวณค่า Money (M) โดยทำการรวมมูลค่าการซื้อของลูกค้าแต่ละรายแสดงดังต่อไปนี้

```
monetary = df_input.groupby('CustomerID')['TotalPrice'].sum().reset_index()
monetary.columns = ['CustomerID', 'Monetary']

rfm_df = rfm_df.merge(monetary, on='CustomerID', how='left')

print("Monetary (M) calculated and merged successfully. Here are the first 5 rows of the RFM DataFrame:")
display(rfm_df.head())
```

Monetary (M) calculated and merged successfully. Here are the first 5 rows of the RFM DataFrame:

	CustomerID	LastPurchaseDate	Recency	Frequency	Monetary
0	12346.0	2024-06-05 10:17:00	325	2	308.88
1	12347.0	2025-04-24 15:52:00	1	7	142230.00
2	12348.0	2025-02-10 13:13:00	74	4	59308.92
3	12349.0	2025-04-08 09:51:00	18	1	57999.15
4	12350.0	2024-06-20 16:01:00	309	1	11035.20

ค่า M (หน่วยเป็นบาท) ของตัวอย่างลูกค้า 5 แถวแรกแสดงดังข้างบน

ตัวอย่าง 9.5 การกำหนดคะแนน RFM และรวมคะแนน

กำหนดคะแนน RFM (เช่น 1-5) ให้กับค่า Recency Frequency และ Monetary สำหรับลูกค้าแต่ละราย โดยใช้การแบ่งกลุ่มแบบ quantile ดังต่อไปนี้

```
# Assign Recency score (higher score for lower recency)
rfm_df['R_Score'] = pd.qcut(rfm_df['Recency'], q=5, labels=[5, 4, 3, 2, 1])

# Assign Frequency score (higher score for higher frequency)
rfm_df['F_Score'] = pd.qcut(rfm_df['Frequency'], q=5, labels=[1, 2, 3, 4, 5])

# Assign Monetary score (higher score for higher monetary value)
rfm_df['M_Score'] = pd.qcut(rfm_df['Monetary'], q=5, labels=[1, 2, 3, 4, 5])
```

```
print("RFM scores assigned successfully. Here are the first 5 rows of the RFM DataFrame with scores:")
display(rfm_df.head())
```

```
rfm_df['R_Score'] = pd.qcut(rfm_df['Recency'].rank(method='first', ascending=True), q=5,
labels=[5, 4, 3, 2, 1]).astype(int)
rfm_df['F_Score'] = pd.qcut(rfm_df['Frequency'].rank(method='first', ascending=True), q=5,
labels=[1, 2, 3, 4, 5]).astype(int)
rfm_df['M_Score'] = pd.qcut(rfm_df['Monetary_x'].rank(method='first', ascending=True), q=5,
labels=[1, 2, 3, 4, 5]).astype(int)

print("RFM scores (R_Score, F_Score, M_Score) calculated successfully using quantile-based
segmentation on ranks. Here are the first 5 rows of the RFM DataFrame with scores:")
display(rfm_df.head())
```

คะแนน RFM ถูกกำหนดและรวมในคอลัมน์สุดท้าย โดยแสดง 5 แถวแรกดังต่อไปนี้

RFM_Score created successfully. Here are the first 5 rows of the RFM DataFrame with the new column:

	CustomerID	LastPurchaseDate	Recency	Frequency	Monetary_x	F_Score	M_Score	R_Score	Monetary_y	RFM_Score
0	12346.0	2024-06-05 10:17:00	325	2	308.88	2	1	1	308.88	121
1	12347.0	2025-04-24 15:52:00	1	7	142230.00	4	5	5	142230.00	545
2	12348.0	2025-02-10 13:13:00	74	4	59308.92	3	4	2	59308.92	234
3	12349.0	2025-04-08 09:51:00	18	1	57999.15	1	4	4	57999.15	414
4	12350.0	2024-06-20 16:01:00	309	1	11035.20	1	2	1	11035.20	112

จะเห็นว่าลูกค้าไอดี 12349 จะมีค่าคะแนน RFM เป็น 414 เพราะ ค่า Recency = 18 (4 คะแนน) ในขณะที่ ค่า Frequency = 1 (1 คะแนน) ทำยอด ค่า Monetary = 57999.15 (4 คะแนน)

ตัวอย่าง 9.6 การกำหนดกลุ่มลูกค้าตามค่าคะแนน RFM

หลังจากที่เราได้คะแนน RFM เราสามารถแบ่งกลุ่มลูกค้าโดยการนิยามลูกค้าแต่ละกลุ่มโดยตรงดังต่อไปนี้

```
def assign_rfm_segment(rfm_score):
    if rfm_score in ['555', '554', '544', '545', '454', '455', '445']:
        return 'Champions'
    elif rfm_score in ['543', '444', '344', '345', '434', '334', '435', '535']:
        return 'Loyal Customers'
    elif rfm_score in ['551', '552', '541', '542', '531', '532', '451', '452', '441', '442', '431', '432']:
        return 'Promising Customers'
    elif rfm_score in ['512', '511', '412', '411', '312', '311']:
        return 'New Customers'
    elif rfm_score in ['222', '223', '233', '322', '323', '333']:
        return 'Customers Needing Attention'
    elif rfm_score in ['122', '123', '133', '112', '111']:
        return 'At Risk'
    elif rfm_score in ['211', '212', '311', '312']:
        return 'Hibernating'
    else:
        return 'Others'

rfm_df['RFM_Segment'] = rfm_df['RFM_Score'].apply(assign_rfm_segment)

print("RFM_Segment created successfully. Here are the first 5 rows of the RFM DataFrame with the new column:")
display(rfm_df.head())
```

RFM_Segment created successfully. Here are the first 5 rows of the RFM DataFrame with the new column:

	CustomerID	LastPurchaseDate	Recency	Frequency	Monetary_x	F_Score	M_Score	R_Score	Monetary_y	RFM_Score	RFM_Segment
0	12346.0	2024-06-05 10:17:00	325	2	308.88	2	1	1	308.88	121	Others
1	12347.0	2025-04-24 15:52:00	1	7	142230.00	4	5	5	142230.00	545	Champions
2	12348.0	2025-02-10 13:13:00	74	4	59308.92	3	4	2	59308.92	234	Others
3	12349.0	2025-04-08 09:51:00	18	1	57999.15	1	4	4	57999.15	414	Others
4	12350.0	2024-06-20 16:01:00	309	1	11035.20	1	2	1	11035.20	112	At Risk

ตัวอย่าง 9.7 การทำ Visualize RFM Segments

ทำการ Visualize การแจกแจงของลูกค้าในแต่ละกลุ่ม แสดงดังต่อไปนี้

```
rfm_segment_counts = rfm_df['RFM_Segment'].value_counts().reset_index()
rfm_segment_counts.columns = ['RFM_Segment', 'Number of Customers']

print("Number of customers in each RFM Segment:")
display(rfm_segment_counts)
```

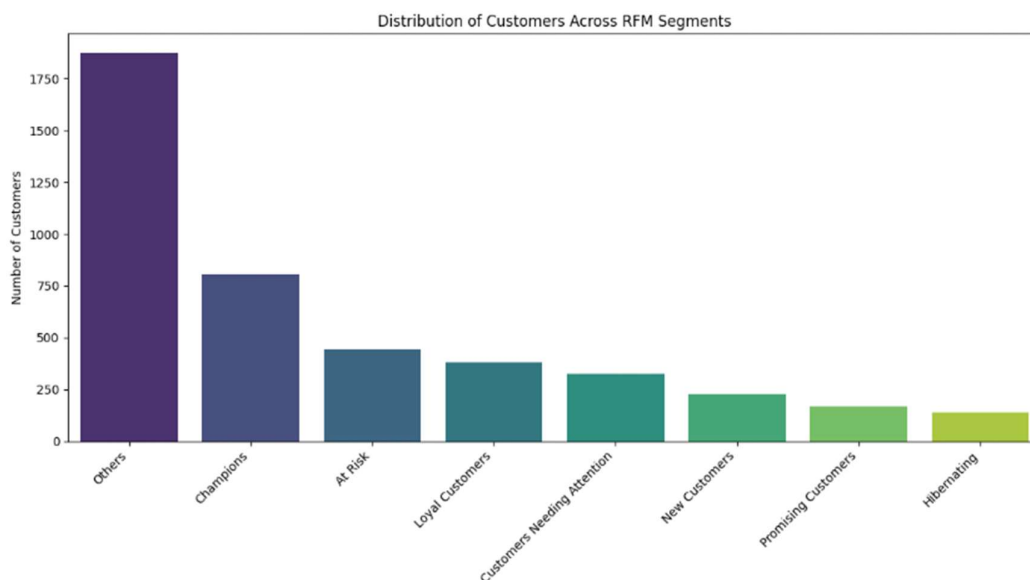
Number of customers in each RFM Segment:

	RFM_Segment	Number of Customers
0	Others	1872
1	Champions	806
2	At Risk	446
3	Loyal Customers	379
4	Customers Needing Attention	326
5	New Customers	230
6	Promising Customers	168
7	Hibernating	141

จากยอดที่นับของจำนวนลูกค้าในแต่ละกลุ่มสามารถนำมาสร้างแผนภูมิแท่งเพื่อช่วยให้เห็นการแจกแจงของลูกค้าในแต่ละกลุ่มดังต่อไปนี้

```
import matplotlib.pyplot as plt
import seaborn as sns

plt.figure(figsize=(12, 7))
sns.barplot(x='RFM_Segment', y='Number of Customers', data=rfm_segment_counts,
palette='viridis', hue='RFM_Segment', legend=False)
plt.title('Distribution of Customers Across RFM Segments')
plt.xlabel('RFM_Segment')
plt.ylabel('Number of Customers')
plt.xticks(rotation=45, ha='right') # Rotate x-axis labels for better readability
plt.tight_layout()
plt.show()
```



ภาพที่ 9.1 การจัดกลุ่ม RFM

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

โดยสรุป การวิเคราะห์ RFM สามารถแบ่งกลุ่มลูกค้าออกเป็นหลายกลุ่มตามค่า Recency, Frequency, และ Monetary ดังนี้

1. **Champions** ลูกค้าที่มีมูลค่าสูงที่สุด (คะแนนสูงทุกมิติ)
2. **Loyal Customers** ลูกค้าประจำที่มีความถี่และการใช้จ่ายสูง
3. **Promising Customers** ลูกค้าที่มีศักยภาพในการเติบโต
4. **New Customers** ลูกค้าใหม่ที่เพิ่งเริ่มซื้อสินค้า
5. **Customers Needing Attention** ลูกค้าที่เริ่มมีพฤติกรรมซื้อลดลง
6. **At Risk** ลูกค้าที่มีความเสี่ยงจะเลิกซื้อ
7. **Hibernating** ลูกค้าที่ไม่เคลื่อนไหว
8. **Others** ลูกค้าที่ไม่เข้ากลุ่มชัดเจน ต้องวิเคราะห์เพิ่มเติม

อย่างไรก็ตามการแบ่งคะแนนด้วย RFM มีปัญหาเนื่องจากค่าซ้ำในตัวแปร Frequency จึงแก้ไขโดยใช้การจัดอันดับก่อน เพื่อให้สามารถแบ่งช่วงข้อมูลได้ถูกต้อง สำหรับการให้คะแนน ถ้า Recency ยิ่งมีค่าน้อย จะยิ่งได้คะแนนสูง แต่สำหรับ Frequency และ Monetary จะกลับกัน ถ้าค่ายิ่งมาก จะยิ่ง

ได้คะแนนสูง หลังจากนั้นคะแนน R, F, และ M จะถูกนำมารวมกันเป็น RFM Score (ตัวอย่างเช่น 555 เป็นต้น) เพื่อใช้ในการจัดกลุ่มลูกค้า

สำหรับการกระจายของกลุ่มลูกค้า พบว่า กลุ่ม Others มีจำนวนมากที่สุด (1,872 คน) รองลงมาคือ Champions (806 คน) At Risk (446 คน)) Loyal Customers (379 คน) Customers Needing Attention (326 คน) New Customers (230 คน) Promising Customers (168 คน) และ Hibernating (141 คน) ตามลำดับ

ตัวอย่าง 9.8 การแบ่งกลุ่มโดยใช้วิธี K-Means Clustering

เพื่อเตรียมข้อมูลสำหรับทำ K-Means Clustering จะทำการสร้าง DataFrame ใหม่โดยการเลือกเฉพาะ “R_Score”, “ F_Score”, และ “M_Score และแสดง 5 แถวแรกดังต่อไปนี้

```
rfm_for_clustering = rfm_df[['R_Score', 'F_Score', 'M_Score']]

print("DataFrame for clustering (rfm_for_clustering) created successfully. Here are its first 5 rows:")

display(rfm_for_clustering.head())
```

DataFrame for clustering (rfm_for_clustering) created successfully. Here are its first 5 rows:

	R_Score	F_Score	M_Score
0	1	2	1
1	5	4	5
2	2	3	4
3	4	1	4
4	1	1	2

ตัวอย่าง 9.9 การหาจำนวน Clusters ที่เหมาะสมที่สุด โดยใช้วิธี Elbow

หาจำนวน Clusters ที่เหมาะสมโดยทำการแปรผันค่า Cluster จาก (1-10) และทำการพล็อตค่า Inertia เพื่อทำการหาค่า K ที่เหมาะสมที่สุด ซึ่งจะเกิดขึ้นที่จุดหักศอกของกราฟ

```
from sklearn.cluster import KMeans
import matplotlib.pyplot as plt

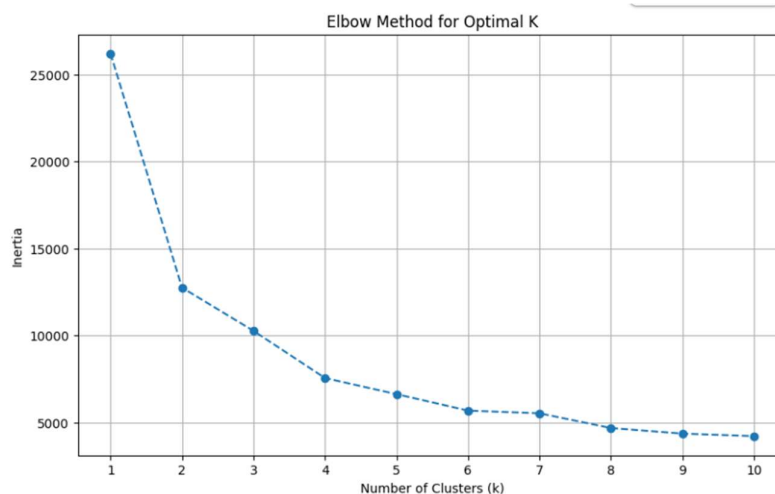
# Initialize an empty list to store inertia values
inertia = []

# Define the range of cluster numbers to test
k_range = range(1, 11) # Testing from 1 to 10 clusters

# Loop through the range of k values, fit KMeans, and store inertia
for k in k_range:
    kmeans = KMeans(n_clusters=k, random_state=42, n_init='auto')
    kmeans.fit(rfm_for_clustering)
    inertia.append(kmeans.inertia_)

# Plot the Elbow Method graph
plt.figure(figsize=(10, 6))
plt.plot(k_range, inertia, marker='o', linestyle='--')
plt.title('Elbow Method for Optimal K')
plt.xlabel('Number of Clusters (k)')
plt.ylabel('Inertia')
plt.xticks(k_range)
plt.grid(True)
plt.show()

print("Elbow Method plot generated successfully.")
```



ภาพที่ 9.2 การหาจำนวน cluster ที่เหมาะสมด้วยวิธี Elbow

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

จากกราฟเราจะเลือกจำนวน Cluster=6 เนื่องจากเป็นจุดหักศอกที่เข้าสู่สภาวะเสถียรมากที่สุด

ตัวอย่าง 9.10 การใช้วิธี K_Means Clustering

ทำการใช้วิธี K_Means Clustering โดยใช้ข้อมูล RFM ในตัวอย่าง 9.8 ดังต่อไปนี้

```
kmeans_model = KMeans(n_clusters=6, random_state=42, n_init='auto')
kmeans_model.fit(rfm_for_clustering)

print("K-Means clustering applied successfully.")
```

หลังจากนั้นทำการกำหนดเลเบลไปที่ RFM DataFrame โดยที่เราจะทำการพยากรณ์ Cluster Labels หลังจากนั้นจะทำการเพิ่มไปยังคอลัมน์สุดท้ายแสดงดังต่อไปนี้

```
rfm_df['RFM_Cluster'] = kmeans_model.predict(rfm_for_clustering)

print("Cluster labels added to rfm_df as 'RFM_Cluster'. Here are the first 5 rows with the new column:")
```

```
display(rfm_df.head())
```

Cluster labels added to rfm_df as 'RFM_Cluster'. Here are the first 5 rows with the new column:

	CustomerID	LastPurchaseDate	Recency	Frequency	Monetary_x	R_Score	F_Score	Monetary_y	M_Score	RFM_Score	RFM_Segment	RFM_Cluster
0	12346.0	2024-06-05 10:17:00	325	2	308.88	1	2	308.88	1	121	Others	3
1	12347.0	2025-04-24 15:52:00	1	7	142230.00	5	4	142230.00	5	545	Champions	4
2	12348.0	2025-02-10 13:13:00	74	4	59308.92	2	3	59308.92	4	234	Others	2
3	12349.0	2025-04-08 09:51:00	18	1	57999.15	4	1	57999.15	4	414	Others	0
4	12350.0	2024-06-20 16:01:00	309	1	11035.20	1	1	11035.20	2	112	At Risk	3

ตัวอย่าง 9.11 การคำนวณค่าเฉลี่ยของคะแนน RFM ต่อ Cluster

ทำการคำนวณค่าเฉลี่ยคะแนน Recency, Frequency, และ Monetary ของ RFM แต่ละ Cluster แสดงดังต่อไปนี้

```
rfm_cluster_means = rfm_df.groupby('RFM_Cluster')[['Recency', 'Frequency',
'Monetary_x']].mean().reset_index()

print("Mean Recency, Frequency, and Monetary scores for each RFM_Cluster:")
display(rfm_cluster_means)
```

Mean Recency, Frequency, and Monetary scores for each RFM_Cluster:

	RFM_Cluster	Recency	Frequency	Monetary_x
0	0	149.981900	1.384615	108494.799122
1	1	12.699248	4.321805	78902.458602
2	2	100.607143	4.637931	99876.498248
3	3	212.553971	1.490835	9010.176843
4	4	13.554307	15.493134	391694.328090
5	5	32.698198	1.569069	12992.345270

ตัวอย่าง 9.12 การทำ Visualize RFM Segments

เริ่มจากการหาจำนวนลูกค้าทั้งหมดของแต่ละกลุ่ม

```
rfm_cluster_distribution = rfm_df['RFM_Cluster'].value_counts().reset_index()
rfm_cluster_distribution.columns = ['RFM_Cluster', 'Number of Customers']

print("Distribution of customers across K-Means clusters:")
display(rfm_cluster_distribution)
```

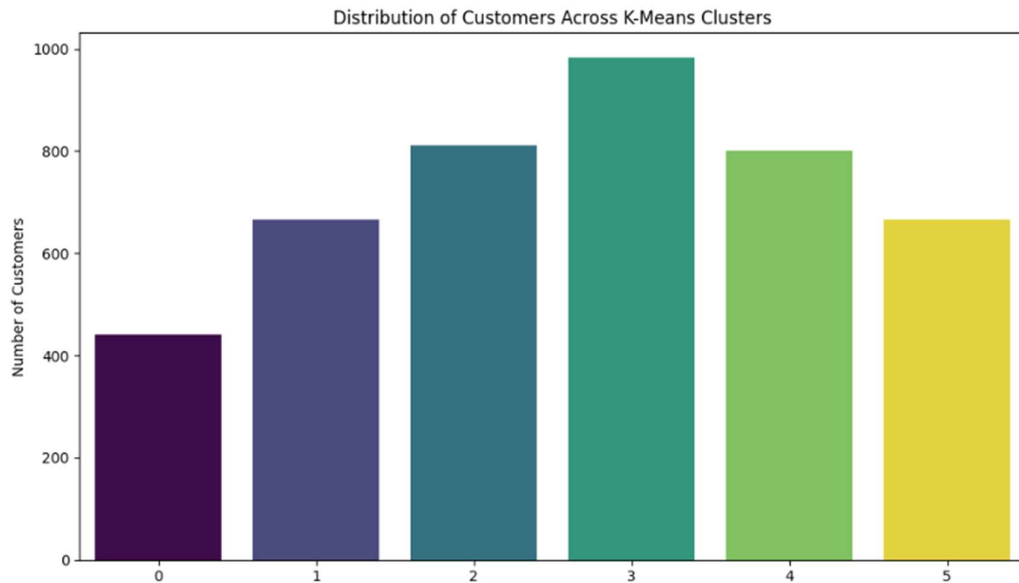
Distribution of customers across K-Means clusters

RFM_Cluster	Number of Customers
3	982
2	812
4	801
5	666
1	665
0	442

ทำการพล็อตกราฟแท่งแสดงดังต่อไปนี้

```
import matplotlib.pyplot as plt
import seaborn as sns

plt.figure(figsize=(10, 6))
sns.barplot(x='RFM_Cluster', y='Number of Customers', data=rfm_cluster_distribution,
            palette='viridis', hue='RFM_Cluster', legend=False)
plt.title('Distribution of Customers Across K-Means Clusters')
plt.xlabel('RFM Cluster')
plt.ylabel('Number of Customers')
plt.xticks(rotation=0)
plt.tight_layout()
plt.show()
```



ภาพที่ 9.3 การแจกแจงจำนวนลูกค้าในแต่ละ cluster

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

โดยสรุปจากการวิเคราะห์ข้อมูลด้วยวิธี K-Means โดยใช้ตัวชี้วัด RFM ได้แก่ Recency (การซื้อล่าสุด) Frequency (ความถี่ในการซื้อ) และ Monetary (มูลค่าการใช้จ่าย) สามารถแบ่งลูกค้าออกเป็น 6 กลุ่มหลัก ซึ่งแต่ละกลุ่มมีพฤติกรรมและคุณค่าทางธุรกิจที่แตกต่างกันอย่างชัดเจน

กลุ่มแรกคือ Cluster 0 (ลูกค้ามียอดซื้อสูงแต่หยุดซื้อ) เป็นกลุ่มลูกค้าที่เคยใช้จ่ายในระดับสูงมาก แต่ไม่ได้ทำการซื้อมาเป็นระยะเวลานาน และมีความถี่ในการซื้อต่ำ แสดงให้เห็นว่าลูกค้ากลุ่มนี้อาจเคยมีความสำคัญต่อธุรกิจ แต่ปัจจุบันเริ่มหยุดซื้อ ดังนั้นจึงควรมีกิจกรรมในการกระตุ้นให้กลับมาซื้ออีกครั้ง

ถัดมาคือ Cluster 1 (ลูกค้าประจำและเพิ่งซื้อ) ซึ่งเป็นกลุ่มลูกค้าที่มีความเคลื่อนไหวสูง มีการซื้อสินค้าอย่างต่อเนื่อง และมีมูลค่าการใช้จ่ายในระดับดี กลุ่มนี้ถือเป็นลูกค้าที่มีความภักดี (Loyal) และมีส่วนสำคัญในการสร้างรายได้อย่างสม่ำเสมอให้กับธุรกิจ

สำหรับ Cluster 2 (ลูกค้าที่มีค่าและยังมีส่วนร่วม) เป็นกลุ่มลูกค้าที่มีมูลค่าการใช้จ่ายสูงและซื้อค่อนข้างบ่อย แต่ระยะเวลาการซื้อครั้งล่าสุดอยู่ในระดับปานกลาง แสดงให้เห็นว่าลูกค้ากลุ่มนี้ยังมีศักยภาพสูง แต่หากไม่ได้รับการดูแลอย่างเหมาะสม อาจลดระดับการมีส่วนร่วมลงได้ในอนาคต

ในขณะที่ Cluster 3 (ลูกค้าที่หยุดซื้อหรือไม่เคลื่อนไหว) เป็นกลุ่มที่มีค่าทุกด้านต่ำ ไม่ว่าจะเป็นความถี่ในการซื้อ มูลค่าการใช้จ่าย หรือความล่าช้าของการซื้อ ลูกค้ากลุ่มนี้มีแนวโน้มที่จะเลิกใช้งานไปแล้ว และอาจต้องใช้ความพยายามสูงในการดึงกลับมา ซึ่งในบางกรณีอาจไม่คุ้มค่าต่อการลงทุน

Cluster 4 (ลูกค้าดีที่สุดหรือ Champions) เป็นกลุ่มที่มีคุณค่าสูงที่สุดในทุกมิติ ทั้งการซื้อบ่อย ซื้อไม่นาน และมีมูลค่าการใช้จ่ายสูงมาก ลูกค้ากลุ่มนี้ถือเป็นหัวใจหลักของธุรกิจ ควรได้รับการดูแลเป็นพิเศษ เช่น การมอบสิทธิพิเศษหรือโปรโมชั่นเฉพาะ เพื่อรักษาความสัมพันธ์ในระยะยาว

สุดท้ายคือ Cluster 5 (ลูกค้าใหม่หรือซื้อครั้งเดียว) ซึ่งเป็นกลุ่มลูกค้าที่เพิ่งเริ่มซื้อสินค้า มีความถี่และมูลค่าการใช้จ่ายยังอยู่ในระดับต่ำ แต่มีความล่าช้าในการซื้อที่ดี กลุ่มนี้มีศักยภาพในการพัฒนาไปเป็นลูกค้าประจำในอนาคต หากได้รับการดูแลและกระตุ้นอย่างเหมาะสม เช่น การทำ Onboarding หรือแคมเปญส่งเสริมการซื้อซ้ำ

โดยสรุป การแบ่งกลุ่มลูกค้าด้วย RFM และ K-Means ช่วยให้ธุรกิจสามารถเข้าใจพฤติกรรมลูกค้าในแต่ละกลุ่มได้อย่างชัดเจน และสามารถวางกลยุทธ์ทางการตลาดให้เหมาะสมกับลูกค้าแต่ละประเภท เพื่อเพิ่มประสิทธิภาพในการรักษาลูกค้าเดิมและสร้างรายได้ในระยะยาวได้อย่างมีประสิทธิภาพ

ตัวอย่าง 9.13 การเปรียบเทียบผลการจัดกลุ่มระหว่าง RFM แบบดั้งเดิม และ แบบผสมกับ K-Means Clustering

การเปรียบเทียบระหว่างวิธีทั้งสองทำได้โดยการใช้ Cross-Tabulation

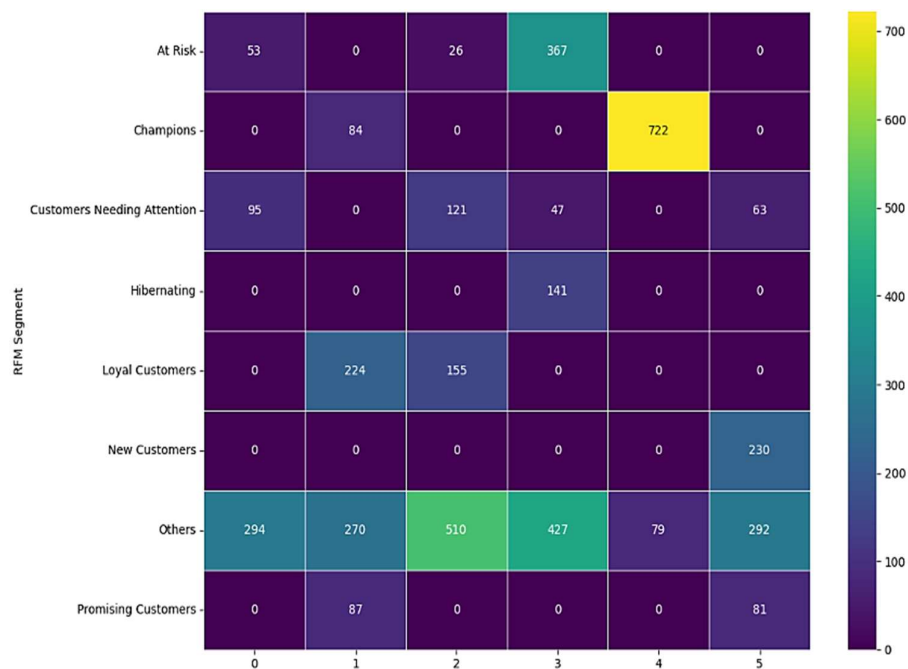
```
crosstab_rfm_kmeans = pd.crosstab(rfm_df['RFM_Segment'], rfm_df['RFM_Cluster'])

print("Cross-tabulation of RFM Segments and K-Means Clusters:")
display(crosstab_rfm_kmeans)
```

Cross-tabulation of RFM Segments and K-Means Clusters:

RFM_Cluster	0	1	2	3	4	5
RFM_Segment						
At Risk	53	0	26	367	0	0
Champions	0	84	0	0	722	0
Customers Needing Attention	95	0	121	47	0	63
Hibernating	0	0	0	141	0	0
Loyal Customers	0	224	155	0	0	0
New Customers	0	0	0	0	0	230
Others	294	270	510	427	79	292
Promising Customers	0	87	0	0	0	81

หลังจากนั้นสามารถ visualize การทับซ้อนของกลุ่มลูกค้าโดยทำการพล็อตกราฟ Heatmap



ภาพที่ 9.4 แผนภาพ Heat Map แสดงการทับซ้อนระหว่าง วิธี RFM กับ K-means clustering

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

แผนภาพ Heatmap นี้แสดงการทับซ้อนระหว่างการแบ่งกลุ่มลูกค้าด้วยวิธี RFM และการจัดกลุ่มด้วยอัลกอริทึม K-Means โดยมีจุดประสงค์เพื่อเปรียบเทียบว่าทั้งสองวิธีให้ผลลัพธ์ที่สอดคล้องกันมากน้อยเพียงใด ภายในภาพ แกนแนวดิ่งแสดงกลุ่มลูกค้าตาม RFM เช่น Champions At Risk Loyal Customers และกลุ่มอื่น ๆ ขณะที่แกนแนวนอนแสดงกลุ่ม K-Means ตั้งแต่ Cluster 0 ถึง Cluster 5 โดยตัวเลขและความเข้มของสีในแต่ละช่องแสดงจำนวนลูกค้าที่อยู่ในทั้งสองกลุ่มพร้อมกัน ซึ่งช่วยให้เห็นการกระจายตัวและความสัมพันธ์ของกลุ่มลูกค้าได้อย่างชัดเจน

จากภาพพบว่าบาง RFM Segment มีความสอดคล้องกับ K-Means อย่างเด่นชัด เช่น กลุ่ม Champions กระจุกตัวอยู่ใน Cluster 4 เป็นหลัก แสดงให้เห็นว่าอัลกอริทึมสามารถแยกกลุ่มลูกค้าที่มีมูลค่าสูงและมีพฤติกรรมดีที่สุดออกมาได้อย่างแม่นยำ เช่นเดียวกับกลุ่ม New Customers ที่อยู่ใน Cluster 5 ทั้งหมด สะท้อนว่าลูกค้าใหม่มีลักษณะเฉพาะที่ชัดเจนจนสามารถถูกจัดอยู่ในกลุ่มเดียวกันได้โดยอัตโนมัติ นอกจากนี้ กลุ่ม Hibernating และ At Risk ยังพบว่ากระจุกตัวใน Cluster 3 ซึ่งบ่งชี้ว่า Cluster นี้เป็นกลุ่มลูกค้าที่มีความเสี่ยงจะเลิกใช้งานหรือไม่มีการเคลื่อนไหว

ในขณะเดียวกัน ยังมีกลุ่มลูกค้าบางประเภทที่ไม่ได้ถูกจัดอยู่ใน Cluster เดียวอย่างชัดเจน เช่น Loyal Customers ที่กระจายอยู่ใน Cluster 1 และ 2 แสดงให้เห็นว่าลูกค้าประจำยังมีความแตกต่างภายใน เช่น ระดับความถี่หรือมูลค่าการใช้จ่ายที่ต่างกัน นอกจากนี้ กลุ่ม Customers Needing Attention และ Promising Customers ก็มีการกระจายอยู่หลาย Cluster ซึ่งสะท้อนถึงความหลากหลายของพฤติกรรมลูกค้าในกลุ่มเหล่านี้ และอาจต้องการการวิเคราะห์เพิ่มเติมเพื่อแบ่งย่อยให้ละเอียดขึ้น

กลุ่ม Others เป็นกลุ่มที่มีการกระจายตัวมากที่สุดในทุก Cluster ซึ่งแสดงถึงความไม่เป็นเนื้อเดียวกันของกลุ่มนี้ โดยอาจเป็นกลุ่มที่รวมลูกค้าที่ไม่เข้าข่าย Segment ใดอย่างชัดเจน จึงถูกกระจายไปตามลักษณะพฤติกรรมจริงที่ K-Means ตรวจจับได้

โดยสรุป แผนภาพนี้แสดงให้เห็นว่า การแบ่งกลุ่มลูกค้าด้วย RFM และ K-Means มีความสอดคล้องกันในหลายกรณี โดยเฉพาะในกลุ่มที่มีลักษณะเด่นชัด เช่น Champions New Customers และลูกค้าที่ไม่ Active ขณะเดียวกันยังช่วยเปิดเผยความซับซ้อนภายในบาง Segment ที่ดูเหมือนเป็นกลุ่มเดียวใน RFM แต่แท้จริงแล้วสามารถแบ่งย่อยได้มากขึ้นด้วยวิธี K-Means ซึ่งเป็นประโยชน์อย่างยิ่งต่อการวางกลยุทธ์ทางการตลาดให้ตรงกับพฤติกรรมของลูกค้าแต่ละกลุ่มอย่างมีประสิทธิภาพ

9.5 สรุปผลลัพธ์จากการจัดกลุ่มลูกค้า

จากกรณีศึกษาการจัดกลุ่มลูกค้าด้วยเทคนิค RFM ร่วมกับการใช้ K-Means Clustering พบว่าสามารถแบ่งลูกค้าออกเป็นกลุ่มที่มีลักษณะพฤติกรรมและมูลค่าทางธุรกิจแตกต่างกันได้อย่างชัดเจน โดยการใช้ข้อมูล Recency Frequency และ Monetary ช่วยสะท้อนพฤติกรรม การซื้อของลูกค้าได้อย่างครอบคลุม และทำให้สามารถวิเคราะห์เชิงลึกได้ดียิ่งขึ้น

ผลการวิเคราะห์ด้วย K-Means สามารถแบ่งลูกค้าออกเป็น 6 กลุ่มหลัก ซึ่งแต่ละกลุ่มมีลักษณะเฉพาะที่แตกต่างกัน โดยกลุ่มแรกคือ Cluster 0 เป็นลูกค้าที่เคยใช้จ่ายสูงแต่หยุดซื้อไปแล้ว มีความถี่ต่ำและไม่ได้ซื้อสินค้าในช่วงเวลาล่าสุด จึงควรมีการกระตุ้นให้กลับมาซื้ออีกครั้ง ถัดมาคือ Cluster 1 ซึ่งเป็นกลุ่มลูกค้าประจำที่ยังมีความเคลื่อนไหวสูง ซื้อสินค้าอย่างสม่ำเสมอ และสร้างรายได้อย่างต่อเนื่อง ถือเป็นกลุ่มลูกค้าสำคัญของธุรกิจ

ในส่วนของ Cluster 2 เป็นกลุ่มลูกค้าที่มีมูลค่าสูงและยังมีส่วนร่วมด้วยธุรกิจในระดับดี แม้ว่าความถี่ในการซื้อจะอยู่ในระดับปานกลาง จึงถือเป็นกลุ่มที่มีศักยภาพและควรได้รับการดูแลอย่างต่อเนื่อง ขณะที่ Cluster 3 เป็นกลุ่มลูกค้าที่ไม่เคลื่อนไหวหรือมีแนวโน้มเลิกซื้อ โดยมีความถี่และมูลค่าการใช้จ่ายต่ำ จึงเป็นกลุ่มที่มีความเสี่ยงสูงและอาจต้องใช้กลยุทธ์เฉพาะในการดึงกลับมา สำหรับ Cluster 4 เป็นกลุ่มลูกค้าที่มีคุณค่าสูงที่สุด หรือที่เรียกว่า Champions ซึ่งมีการซื้อบ่อย ใช้จ่ายสูง และมีความเคลื่อนไหวล่าสุด ถือเป็นกลุ่มหลักที่สร้างรายได้ให้กับธุรกิจและควรได้รับการดูแลเป็นพิเศษ ส่วน Cluster 5 เป็นกลุ่มลูกค้าใหม่หรือซื้อเพียงครั้งเดียว ซึ่งยังมีมูลค่าและความถี่ต่ำ แต่มีศักยภาพในการพัฒนาเป็นลูกค้าประจำในอนาคต

นอกจากนี้ เมื่อเปรียบเทียบผลการจัดกลุ่มระหว่าง RFM และ K-Means พบว่าทั้งสองวิธีมีความสอดคล้องกันในหลายกลุ่ม เช่น กลุ่ม Champions และ New Customers ที่ถูกจัดอยู่ใน cluster เดียวกันอย่างชัดเจน ขณะที่บางกลุ่ม เช่น Loyal Customers หรือ Customers Needing Attention มีการกระจายอยู่หลาย cluster แสดงถึงความหลากหลายภายในกลุ่มลูกค้า โดยสรุปการใช้ RFM ร่วมกับ K-Means ช่วยให้องค์กรสามารถเข้าใจพฤติกรรมลูกค้าได้อย่างลึกซึ้งมากขึ้น สามารถระบุลูกค้ากลุ่มสำคัญ กลุ่มที่มีความเสี่ยง และกลุ่มที่มีศักยภาพในการเติบโต ซึ่งเป็นประโยชน์อย่างยิ่งต่อการวางแผนกลยุทธ์ทางการตลาด การรักษาลูกค้า และการเพิ่มประสิทธิภาพในการดำเนินธุรกิจในระยะยาว

จากกรณีศึกษาการจัดกลุ่มลูกค้าด้วยเทคนิค RFM ร่วมกับการใช้ K-Means Clustering พบว่าสามารถแบ่งลูกค้าออกเป็น 6 กลุ่มหลักที่สะท้อนพฤติกรรมและมูลค่าทางธุรกิจได้อย่างชัดเจน โดยใช้ตัวแปร Recency, Frequency และ Monetary เป็นตัวแทนพฤติกรรม การซื้อของลูกค้า ผลการวิเคราะห์แสดงให้เห็นว่า Cluster 4 (Champions) เป็นกลุ่มลูกค้าที่มีคุณค่าสูงที่สุด ทั้งในด้านความถี่ในการซื้อ ความล่าสุด และมูลค่าการใช้จ่าย ขณะที่ Cluster 1 และ Cluster 2 เป็นกลุ่มลูกค้าประจำ

และลูกค้าที่มีมูลค่าสูง ซึ่งยังมีความสัมพันธ์กับธุรกิจในระดับดี ส่วน Cluster 3 เป็นกลุ่มลูกค้าที่มีความเสี่ยงจะเลิกซื้อหรือไม่เคลื่อนไหว และ Cluster 0 เป็นลูกค้าที่เคยใช้จ่ายสูงแต่หยุดซื้อไปแล้ว ขณะที่ Cluster 5 เป็นกลุ่มลูกค้าใหม่ที่ยังมีโอกาสเติบโตในอนาคต

เมื่อพิจารณาการเปรียบเทียบระหว่าง RFM และ K-Means พบว่าหลายกลุ่มมีความสอดคล้องกันอย่างชัดเจน เช่น Champions, New Customers และกลุ่มลูกค้าที่ไม่ Active ซึ่งสะท้อนว่าโมเดลสามารถจับพฤติกรรมลูกค้าได้แม่นยำ อย่างไรก็ตาม จุดที่น่าสนใจอย่างมากคือ กลุ่ม “Others” ซึ่งใน RFM ถูกจัดเป็นกลุ่มที่ไม่เข้าข่ายชัดเจน กลับถูก K-Means แยกออกเป็นหลาย Cluster อย่างชัดเจน ลูกค้าในกลุ่ม Others ไม่ได้มีลักษณะเหมือนกันทั้งหมด แต่มีความหลากหลายสูง (Heterogeneous) โดย K-Means สามารถแยกย่อยลูกค้ากลุ่มนี้ออกไปอยู่ในหลาย Cluster บางส่วนอยู่ใน Cluster ที่มีลักษณะใกล้เคียงลูกค้าประจำ (Cluster 1 และ 2) ขณะที่บางส่วนอยู่ใน Cluster ที่มีพฤติกรรมใกล้เคียงลูกค้าที่เริ่มหายไป (Cluster 3) หรือแม้แต่กลุ่มลูกค้าใหม่ (Cluster 5) ผลลัพธ์นี้สะท้อนให้เห็นว่า การแบ่งกลุ่มแบบ RFM เพียงอย่างเดียวอาจยังไม่สามารถจับ “ความแตกต่างภายในกลุ่ม Others” ได้อย่างละเอียด เนื่องจากการแบ่งตามกฎที่กำหนดไว้ล่วงหน้า แต่ K-Means ซึ่งเป็นการเรียนรู้จากข้อมูลจริง สามารถค้นหารูปแบบที่ซ่อนอยู่และแยกกลุ่มลูกค้าออกตามพฤติกรรมที่แท้จริงได้ดีกว่า ดังนั้น กลุ่ม Others จึงถือเป็นกลุ่มที่มีศักยภาพในการวิเคราะห์เพิ่มเติม เพราะอาจประกอบไปด้วยลูกค้าหลายประเภท เช่น ลูกค้าที่กำลังจะกลายเป็นลูกค้าประจำ ลูกค้าที่เริ่มลดการซื้อ หรือแม้แต่ลูกค้าที่มีมูลค่าสูงแต่พฤติกรรมไม่สม่ำเสมอ การเข้าใจความหลากหลายภายในกลุ่มนี้จะช่วยให้องค์กรสามารถออกแบบกลยุทธ์เฉพาะเจาะจงได้ดียิ่งขึ้น

โดยสรุป การใช้ K-Means ร่วมกับ RFM ไม่เพียงช่วยยืนยันโครงสร้างของกลุ่มลูกค้าเดิม แต่ยังช่วยเปิดเผย “ความซับซ้อนภายในกลุ่มที่ไม่ชัดเจน” โดยเฉพาะกลุ่ม Others ซึ่งเป็นประโยชน์อย่างยิ่งต่อการพัฒนา customer segmentation ให้มีความแม่นยำและนำไปใช้เชิงกลยุทธ์ได้อย่างมีประสิทธิภาพมากขึ้นในระยะยาว

เอกสารอ้างอิง

- Chen, D., Sain, S. L., & Guo, K. (2012). Data mining for the online retail industry: A case study of RFM model-based customer segmentation. *Journal of Database Marketing & Customer Strategy Management*, 19(3), 197-208. <https://doi.org/10.1057/dbm.2012.17>.
- Ngai, E. W. T., Xiu, L., & Chau, D. C. K. (2009). Application of data mining techniques in customer relationship management: A literature review and classification. *Expert Systems with Applications*, 36(2), 2592-2602. <https://doi.org/10.1016/j.eswa.2008.02.021>.
- Jain, A. K. (2010). Data clustering: 50 years beyond K-means. *Pattern Recognition Letters*, 31(8), 651-666. <https://doi.org/10.1016/j.patrec.2009.09.011>.
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77-84. <https://doi.org/10.1111/jbl.12010>.
- Guo, X., Gao, L., Liu, X., & Yin, J. (2017). Improved deep embedded clustering with local structure preservation. In *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence (IJCAI)* (pp. 1753-1759). <https://doi.org/10.24963/ijcai.2017/243>.

บทที่ 10

จากทฤษฎีสู่การปฏิบัติ: การพยากรณ์อุปสงค์ด้วยการเรียนรู้ของเครื่อง

From Theory to Practice: Machine Learning for Demand Forecasting

หัวข้อ

- 10.1 วัตถุประสงค์ของกรณีศึกษา
- 10.2 บริบทธุรกิจ
- 10.3 ชุดข้อมูลและตัวแปรที่ใช้
- 10.4 ขั้นตอนการวิเคราะห์ด้วย การเรียนรู้ของเครื่อง
- 10.5 สรุปผลลัพธ์จากการพยากรณ์อุปสงค์

การพยากรณ์อุปสงค์ (demand forecasting) เป็นกระบวนการสำคัญในการคาดการณ์ความต้องการของสินค้า หรือบริการในอนาคต โดยอาศัยข้อมูลในอดีตและเทคนิคทางคณิตศาสตร์ หรือสถิติ เพื่อให้ได้ผลการพยากรณ์ที่มีความแม่นยำ ซึ่งมีบทบาทอย่างยิ่งต่อการวางแผนในซัพพลายเชน ไม่ว่าจะเป็นการวางแผนการผลิต การจัดสรรทรัพยากร เช่น เครื่องจักร แรงงาน และการจัดซื้อวัตถุดิบ นอกจากนี้ ความแม่นยำของการพยากรณ์ยังส่งผลโดยตรงต่อประสิทธิภาพของระบบการจัดการสินค้าคงคลัง หากการพยากรณ์คลาดเคลื่อน อาจนำไปสู่ปัญหาสินค้าคงคลังส่วนเกิน ซึ่งก่อให้เกิดต้นทุนการจัดเก็บและการเสื่อมสภาพของสินค้า หรือในทางกลับกัน อาจเกิดภาวะสินค้าขาดแคลน ส่งผลให้สูญเสียโอกาสทางการขายและความพึงพอใจของลูกค้า (Waller & Fawcett, 2013)

ในบริบทของการเปิดตัวสินค้าใหม่ โดยเฉพาะในอุตสาหกรรมที่มีการแข่งขันสูง เช่น โทรศัพท์มือถือ ความแม่นยำของการพยากรณ์อุปสงค์ก็มีความสำคัญมากขึ้น เนื่องจากความไม่แน่นอนของความต้องการและพฤติกรรมผู้บริโภค จากกรณีศึกษาพบว่า การพยากรณ์ที่ไม่แม่นยำทำให้เกิดปัญหาสินค้าขาดตลาด ลูกค้าต้องรอสินค้านานถึง 8-10 สัปดาห์ และบางส่วนเปลี่ยนไปซื้อสินค้าของคู่แข่ง ส่งผลให้สูญเสียรายได้และลูกค้า ดังนั้น จึงมีความจำเป็นที่จะต้องพัฒนาวิธีการพยากรณ์อุปสงค์ที่เหมาะสมสำหรับสินค้าใหม่ โดยนำเทคนิคการวิเคราะห์อนุกรมเวลาและการเรียนรู้ของเครื่อง (ML) มาใช้เปรียบเทียบประสิทธิภาพของแบบจำลองต่าง ๆ เช่น ARIMA Exponential Smoothing

Holt-Winters และโครงข่ายประสาทเทียม (Artificial Neural Network, ANN) เพื่อเลือกวิธีที่ให้ ความแม่นยำสูงที่สุด และสามารถนำไปประยุกต์ใช้ในการวางแผนการผลิตและการจัดการสินค้าคง คลังได้อย่างมีประสิทธิภาพมากยิ่งขึ้น (Areerakulkan et al., 2024)

โดยงานวิจัยของ Areerakulkan et al. (2024) มุ่งเน้นการพัฒนาความแม่นยำของการ พยากรณ์อุปสงค์สำหรับสินค้าใหม่ โดยใช้กรณีศึกษาผลิตภัณฑ์โทรศัพท์มือถือ ซึ่งประสบปัญหาการ พยากรณ์ที่คลาดเคลื่อน ส่งผลให้เกิดทั้งภาวะสินค้าขาดแคลนและการวางแผนการผลิตที่ไม่มี ประสิทธิภาพ ผู้วิจัยได้วิเคราะห์ข้อมูลยอดขายในอดีตซึ่งมีลักษณะเป็นอนุกรมเวลาและมีทั้งแนวโน้ม และฤดูกาล พร้อมทั้งเปรียบเทียบวิธีการพยากรณ์ทั้งแบบดั้งเดิม ได้แก่ Holt-Winters ARIMA และ Exponential Smoothing (ETS) กับวิธีการเรียนรู้ของเครื่องแบบโครงข่ายประสาทเทียม (ANN) เพื่อ คัดเลือกวิธีที่เหมาะสมที่สุด ผลการศึกษาพบว่าแบบจำลอง ANN ให้ความแม่นยำสูงที่สุดเมื่อพิจารณา จากตัวชี้วัด เช่น MAD RMSE และ MAPE โดยมีความแม่นยำมากกว่าวิธีเดิมถึง 51.28% นอกจากนี้ ยังมีการนำผลการพยากรณ์ไปเชื่อมโยงกับการวางแผนการผลิตและการบริหารสินค้าคงคลัง ทั้งใน ส่วนของสินค้าสำเร็จรูปและงานระหว่างผลิต เพื่อกำหนดปริมาณการผลิตและระดับสินค้าคงคลังที่ เหมาะสม ส่งผลให้สามารถลดปัญหาการขาดสินค้า ลดการสูญเสียยอดขาย และลดต้นทุนรวมของ ระบบได้อย่างมีนัยสำคัญ โดยต้นทุนลดลงประมาณ 27.71% ซึ่งสะท้อนให้เห็นถึงประสิทธิภาพของ เทคนิคการเรียนรู้ของเครื่องในการยกระดับการพยากรณ์อุปสงค์และการบริหารจัดการซัพพลายเชน ได้อย่างเป็นรูปธรรม

อย่างไรก็ตามในช่วงหลายปีที่ผ่านมา นอกเหนือจากการใช้ ANN เทคนิค ML ได้ถูกนำมาใช้ ในการพยากรณ์อุปสงค์มากขึ้น โดยเฉพาะแบบจำลองประเภท ensemble learning ซึ่งสามารถรวม ผลลัพธ์จากหลายแบบจำลองเพื่อเพิ่มความแม่นยำในการพยากรณ์ หนึ่งในวิธีที่ได้รับความนิยมคือ Random Forest ซึ่งพัฒนาโดย Breiman (2001) โดยใช้แนวคิดของการสร้างต้นไม้ตัดสินใจหลาย ต้นและรวมผลลัพธ์เพื่อลดความแปรปรวนของโมเดล ทำให้สามารถจัดการกับข้อมูลที่มีความซับซ้อน และลดปัญหา overfitting ได้อย่างมีประสิทธิภาพ นอกจากนี้ XGBoost (Extreme Gradient Boosting) ซึ่งพัฒนาโดย Chen และ Guestrin (2016) เป็นอีกหนึ่งแบบจำลองที่ได้รับความนิยม อย่างสูงในงานพยากรณ์อุปสงค์ เนื่องจากใช้แนวคิดของ boosting ในการสร้างแบบจำลองอย่าง ต่อเนื่อง โดยแต่ละรอบจะมุ่งเน้นแก้ไขข้อผิดพลาดของรอบก่อนหน้า ส่งผลให้ได้โมเดลที่มีความ แม่นยำสูงและสามารถจัดการกับข้อมูลขนาดใหญ่ได้อย่างมีประสิทธิภาพ อีกทั้งยังรองรับการจัดการ ข้อมูลที่ขาดหาย (missing data) และสามารถควบคุมความซับซ้อนของโมเดลได้ดี งานวิจัยจำนวน มากได้แสดงให้เห็นว่า Random Forest และ XGBoost มีประสิทธิภาพสูงในการพยากรณ์อุปสงค์ โดยเฉพาะในกรณีที่ข้อมูลมีความไม่เป็นเชิงเส้นหรือมีปัจจัยหลายมิติ เช่น โปรโมชัน ราคา และฤดูกาล Carbonneau et al. (2008) พบว่าเทคนิค ML สามารถให้ผลลัพธ์ที่แม่นยำกว่าวิธีการทางสถิติแบบ

ดั้งเดิม ขณะที่ Makridakis et al. (2018) ชี้ให้เห็นว่า ensemble และ ML มีศักยภาพสูงในการพยากรณ์ข้อมูลอนุกรมเวลาในระดับอุตสาหกรรม

แม้ว่า Random Forest และ XGBoost จะมีข้อดีด้านความแม่นยำและความสามารถในการจัดการข้อมูลที่ซับซ้อน แต่ก็ยังมีข้อจำกัด เช่น ความยากในการตีความผลลัพธ์เมื่อเทียบกับแบบจำลองเชิงเส้น และความจำเป็นในการปรับแต่งพารามิเตอร์ (hyperparameter tuning) อย่างเหมาะสม อย่างไรก็ตาม ด้วยความสามารถในการเรียนรู้รูปแบบข้อมูลที่หลากหลายและความยืดหยุ่นในการประยุกต์ใช้ ทำให้ทั้งสองวิธีเป็นเครื่องมือที่มีประสิทธิภาพสูงสำหรับการพยากรณ์อุปสงค์ในยุคของข้อมูลขนาดใหญ่

10.1 วัตถุประสงค์ของกรณีศึกษา

วัตถุประสงค์ของกรณีศึกษานี้มุ่งเน้นการพัฒนาแบบจำลองการพยากรณ์อุปสงค์โดยใช้เทคนิคการเรียนรู้ของเครื่อง เพื่อเพิ่มความแม่นยำในการคาดการณ์ยอดขายสินค้า โดยอาศัยข้อมูลยอดขายรายสัปดาห์ร่วมกับตัวแปรที่เกี่ยวข้องหลายมิติ เช่น ราคา โปรโมชัน และคุณลักษณะของผู้บริโภค ซึ่งเป็นข้อมูลแบบอนุกรมเวลาหลายตัวแปร (multivariate time series) ทั้งนี้ กรณีศึกษานี้ให้ความสำคัญกับการประยุกต์ใช้เทคนิค Random Forest และ XGBoost ในการสร้างแบบจำลองพยากรณ์ เนื่องจากทั้งสองวิธีมีความสามารถในการจัดการข้อมูลที่ซับซ้อนและความสัมพันธ์แบบไม่เป็นเชิงเส้นได้อย่างมีประสิทธิภาพ

นอกจากนี้ กรณีศึกษายังมีวัตถุประสงค์เพื่อเปรียบเทียบประสิทธิภาพของแบบจำลองทั้งสองโดยใช้ตัวชี้วัดความแม่นยำ เช่น mean squared error (MSE) ค่า R-squared (R^2) และ root mean square percentage error (RMSPE) เพื่อประเมินความเหมาะสมของแต่ละเทคนิคในการพยากรณ์อุปสงค์ อีกทั้งยังมุ่งศึกษาความสำคัญของปัจจัยต่าง ๆ ที่มีผลต่อยอดขายสินค้า ผ่านการวิเคราะห์ feature importance เพื่อทำความเข้าใจว่าปัจจัยใดมีอิทธิพลต่อความต้องการสินค้า นอกจากนี้ กรณีศึกษานี้มีเป้าหมายเพื่อแสดงให้เห็นถึงศักยภาพของ ML ในการพยากรณ์อุปสงค์ และสนับสนุนการตัดสินใจในงานบริหารซัพพลายเชน เช่น การวางแผนการผลิต การบริหารสินค้าคงคลัง และการจัดการด้านการตลาด โดยอาศัยข้อมูลและแบบจำลองเชิงวิเคราะห์ที่มีความแม่นยำและเชื่อถือได้

10.2 บริบทธุรกิจ

กรณีศึกษาอยู่ในบริบทของธุรกิจค้าปลีกสินค้าอุปโภคบริโภค โดยมุ่งเน้นการวิเคราะห์และพยากรณ์ยอดขายของน้ำส้มบรรจุขวดขนาด 64 ออนซ์ ซึ่งจำหน่ายผ่านร้านค้าจำนวน 83 แห่งในพื้นที่รัฐชิคาโก ประเทศสหรัฐอเมริกา โดยสินค้าในตลาดประกอบด้วย 3 แบรินด์หลัก ได้แก่ Dominicks Minute Maid และ Tropicana ซึ่งมีการแข่งขันกันทั้งด้านราคา การส่งเสริมการขาย และการเข้าถึงลูกค้า

ลักษณะของธุรกิจมีความซับซ้อนเนื่องจากยอดขายไม่ได้ขึ้นอยู่กับปัจจัยเดียว แต่ได้รับอิทธิพลจากหลายปัจจัยร่วมกัน เช่น ราคา โปรโมชัน (การโฆษณา) ช่วงเวลา (สัปดาห์) รวมถึงลักษณะทางประชากรศาสตร์ของลูกค้า เช่น อายุ ระดับการศึกษา รายได้ และโครงสร้างครัวเรือน ซึ่งทำให้พฤติกรรมการณ์ซื้อมีความหลากหลายและเปลี่ยนแปลงตามบริบทของตลาด นอกจากนี้ ธุรกิจยังต้องเผชิญกับความท้าทายในการบริหารจัดการอุปสงค์ให้สอดคล้องกับอุปทาน เนื่องจากหากพยากรณ์ต่ำกว่าความเป็นจริงอาจทำให้เกิดปัญหาสินค้าขาดสต็อก ในขณะที่การพยากรณ์สูงเกินไปจะส่งผลให้เกิดต้นทุนสินค้าคงคลังที่สูง ดังนั้น การนำเทคนิคการเรียนรู้ของเครื่อง เช่น Random Forest และ XGBoost เข้ามาช่วยวิเคราะห์ข้อมูลแบบหลายตัวแปร (multivariate time series) จึงเป็นแนวทางสำคัญในการเพิ่มความแม่นยำของการพยากรณ์ และสนับสนุนการตัดสินใจด้านการจัดการซัพพลายเชนของธุรกิจได้อย่างมีประสิทธิภาพ

10.3 ชุดข้อมูลและตัวแปรที่ใช้

กรณีศึกษาใช้ข้อมูลยอดขายสินค้าในรูปแบบอนุกรมเวลา โดยเป็นข้อมูลยอดขายรายสัปดาห์ของน้ำส้มบรรจุขวดขนาด 64 ออนซ์ จากร้านค้าจำนวน 83 แห่งในพื้นที่รัฐชิคาโก ประเทศสหรัฐอเมริกา ซึ่งครอบคลุมสินค้าจำนวน 3 แบรินด์ ได้แก่ Dominicks Minute Maid และ Tropicana โดยข้อมูลมีลักษณะเป็นข้อมูลหลายตัวแปร (Multivariate Data) และถูกจัดเก็บในรูปแบบแถว รวมทั้งสิ้น 28,947 รายการ ตัวแปรที่ใช้ในการวิเคราะห์สามารถแบ่งออกเป็นตัวแปรตามและตัวแปรอิสระ โดยตัวแปรตามคือ LOGMOVE ซึ่งเป็นค่าลอการิทึมของยอดขายสินค้า ใช้แทนปริมาณอุปสงค์สินค้า ขณะที่ตัวแปรอิสระประกอบด้วยหลายมิติ ได้แก่ ด้านสินค้าและการตลาด เช่น PRICE (ราคา) FEAT (การมีหรือไม่มีโฆษณา/โปรโมชัน) BRAND (ยี่ห้อสินค้า) และ WEEK (ช่วงเวลา) ซึ่งสะท้อนพฤติกรรมการณ์ซื้อในแต่ละช่วงเวลา นอกจากนี้ ยังมีตัวแปรด้านประชากรศาสตร์ของพื้นที่ร้านค้า เช่น AGE60 (สัดส่วนประชากรอายุ 60 ปีขึ้นไป) EDUC (ระดับการศึกษา) ETHNIC

(เชื้อชาติ) INCOME (รายได้) HHLARGE (ขนาดครัวเรือน) WORKWOM (สัดส่วนผู้หญิงทำงาน) และ HVAL150 (มูลค่าที่อยู่อาศัย) ซึ่งช่วยอธิบายความแตกต่างของพฤติกรรมผู้บริโภคในแต่ละพื้นที่ ในส่วนของตัวแปรด้านร้านค้าและการแข่งขัน ประกอบด้วย STORE (รหัสร้านค้า) SSTRDIST (ระยะทางถึงร้านค้าใกล้เคียง) SSTRVOL (อัตราส่วนยอดขายของร้าน) CPDIST5 (ระยะทางเฉลี่ยไปยังร้านคู่แข่ง 5 แห่ง) และ CPWVOL5 (อัตราส่วนยอดขายเทียบกับคู่แข่ง) ซึ่งช่วยสะท้อนอิทธิพลของทำเลที่ตั้งและการแข่งขันในตลาด

การใช้ข้อมูลและตัวแปรที่ครอบคลุมหลายมิติดังกล่าว ทำให้สามารถนำไปประยุกต์ใช้กับแบบจำลองการเรียนรู้ของเครื่อง เช่น Random Forest และ XGBoost เพื่อเรียนรู้ความสัมพันธ์ที่ซับซ้อนระหว่างปัจจัยต่าง ๆ และเพิ่มความแม่นยำในการพยากรณ์อุปสงค์สินค้าได้อย่างมีประสิทธิภาพ สรุปตัวแปรที่เกี่ยวข้องทั้งหมดแสดงดังต่อไปนี้

STORE:	หมายเลขร้าน
BRAND:	ยี่ห้อน้ำส้ม (Dominicks Minute Maid Tropicana)
WEEK:	หมายเลขสัปดาห์
LOGMOVE:	ค่า log ของยอดขาย
PRICE:	ราคา
FEAT:	การโฆษณาสินค้า (1, 0)
AGE60:	สัดส่วนของประชากรที่มีอายุตั้งแต่ 60 ปีขึ้นไป
EDUC:	สัดส่วนของประชากรที่จบมหาวิทยาลัย
ETHNIC:	สัดส่วนของประชากรที่ผิวดำหรือละติน
INCOME:	ค่า log ค่ากลางของรายได้
HHLARGE:	สัดส่วนของครัวเรือนที่มีสมาชิกมากกว่า 5 คนขึ้นไป
WORKWOM:	สัดส่วนของผู้หญิงทำงานเต็มเวลา
HVAL150:	สัดส่วนของครัวเรือนที่มีมูลค่ามากกว่า \$150,000
SSTRDIST:	ระยะทางถึงสโตร์ที่ใกล้ที่สุด
SSTRVOL:	อัตราส่วนของยอดขายของสโตร์นี้กับคลังสินค้าที่ใกล้ที่สุด
CPDIST5:	ค่าเฉลี่ยระยะทางในหน่วยไมล์ไปยังซูเปอร์มาร์เก็ตที่ใกล้ที่สุด 5 แห่ง
CPWVOL5:	อัตราส่วนของยอดขายของสโตร์นี้กับค่าเฉลี่ยยอดขายของซูเปอร์มาร์เก็ต 5 แห่ง

(หมายเหตุ: สามารถดาวน์โหลดข้อมูลได้ที่: <https://github.com/gchoi/Dataset/blob/master/oj.csv>)

10.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง

ในหัวข้อนี้จะใช้วิธีการเรียนรู้ของเครื่อง 2 วิธีที่นิยมใช้ ได้แก่ วิธี Random Forest (RF) และวิธี XGboost โดยจะอธิบายประกอบกับการใช้ PYTHON ช่วยในการวิเคราะห์ข้อมูล

10.4.1 Random Forest (RF)

RF เป็นอัลกอริทึมที่สร้างมาจากแนวคิดของ ensemble learning โดยการรวมการตัดสินใจจากหลาย ๆ ต้นไม้ตัดสินใจ เพื่อให้ได้ผลลัพธ์ที่แม่นยำและมีเสถียรภาพมากขึ้น โดยมีขั้นตอนการทำงานดังต่อไปนี้

1. การแบ่งข้อมูลออกเป็นชุด training และ testing ซึ่งในบางกรณีอาจมีความจำเป็นที่จะต้องปรับสเกลของข้อมูล (normalization/scaling)
2. การสุ่มตัวอย่างข้อมูลจากชุด training ด้วยวิธี bootstrap sampling เพื่อสร้าง subset ของข้อมูลจำนวนมาก จากการกำหนดค่าพารามิเตอร์ ($n_estimators$)
3. การสร้างต้นไม้ตัดสินใจโดยจะสร้างจาก subset ที่ได้จากขั้นตอนที่ 2 โดย 1 subset จะถูกสร้างเป็น ต้นไม้ตัดสินใจหนึ่งต้น โดยจะสร้างตามเงื่อนไขการหยุดที่กำหนดตามพารามิเตอร์ อย่างเช่น ความลึกสูงสุดของต้นไม้ (max_depth) โดยไม่เกินข้อมูลโหนดที่กำหนดโดยค่าพารามิเตอร์ (min_sample_leaf) นอกจากนั้น ในแต่ละโหนดของต้นไม้ จะสุ่มเลือกตัวแปร จำนวนหนึ่งเพื่อใช้แทนตัวแปรทั้งหมด
4. การรวมผลลัพธ์จาก decision tree หลายๆต้น ด้วยการใช้ classification และ regression
5. ทำการประเมินผลด้วยชุด testing โดยใช้เมตริกต่างๆเช่น accuracy precision recall (สำหรับ classification) หรือ RMSE MAE (สำหรับ regression)
6. การปรับแต่งพารามิเตอร์เพื่อเพิ่มประสิทธิภาพโดยใช้เทคนิค grid search หรือ random search ร่วมกับ cross validation เพื่อหาชุดพารามิเตอร์ที่เหมาะสมที่สุด

ตัวอย่างที่ 10.1 แสดงขั้นตอนการใช้เทคนิค Random Forest พยากรณ์อุปสงค์ของชุดข้อมูลกรณีศึกษา ด้วย PYTHON ในขั้นตอนที่จำเป็นและแสดงโค้ดประกอบการอธิบายดังต่อไปนี้

ตัวอย่าง 10.1 การ import packages ที่ต้องใช้และข้อมูลนำเข้าเพื่อเริ่มทำการวิเคราะห์

การ import package ที่จำเป็นต้องใช้:

```
import warnings
warnings.filterwarnings("ignore")

# Import packages
import numpy as np
import pandas as pd
import datetime as dt

# Data visualisation
import matplotlib.pyplot as plt
import seaborn as sns; sns.set()
%matplotlib inline

# Time series analysis
from sklearn.ensemble import RandomForestRegressor
import xgboost as xgb
```

หลังจากนั้นสามารถโหลดไฟล์ dataset ที่ต้องใช้ในที่นี้คือไฟล์ชื่อ “oj_ready.csv” และทำการตรวจสอบข้อมูลเบื้องต้น

```
# read csv file name oj.csv
data = pd.read_csv('oj_ready.csv')
data.head()
data.describe()
data.info()
```

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 28947 entries, 0 to 28946
Data columns (total 18 columns):
#   Column      Non-Null Count  Dtype
---  ---
0    store      28947 non-null  int64
1    brand       28947 non-null  object
2    brandno     28947 non-null  int64
3    week        28947 non-null  int64
4    logmove     28947 non-null  float64
5    feat       28947 non-null  int64
6    price      28947 non-null  float64
7    AGE60      28947 non-null  float64
8    EDUC       28947 non-null  float64
9    ETHNIC     28947 non-null  float64
10   INCOME     28947 non-null  float64
11   HHLARGE    28947 non-null  float64
12   WORKWOM    28947 non-null  float64
13   HVAL150    28947 non-null  float64
14   SSTRDIST   28947 non-null  float64
15   SSTRVOL    28947 non-null  float64
16   CPDIST5    28947 non-null  float64
17   CPWVOL5    28947 non-null  float64
dtypes: float64(13), int64(4), object(1)
memory usage: 4.0+ MB
```

ตัวอย่าง 10.2 การแบ่งข้อมูลเป็นชุด Training (80%) และ Testing set (20%)

```
# divide to train and test set
from sklearn.model_selection import train_test_split

# your target variable is 'logmove'
# and you want to remove columns 'brand'

X = data.drop(['brand', 'logmove'], axis=1)
y = data['logmove']

# Split the data into training and testing sets
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)

# Print the shapes of the resulting sets
print("X_train shape:", X_train.shape)
print("X_test shape:", X_test.shape)
print("y_train shape:", y_train.shape)
print("y_test shape:", y_test.shape)
y_train
```

ตัวอย่าง 10.3 Fit โมเดลโดยใช้ Random Forest Regressor Model (random state = 42)

```
# random forest

# Create a Random Forest Regressor model
rf_model = RandomForestRegressor(random_state=42)

# Train the model on the training data
rf_model.fit(X_train, y_train)

# Make predictions on the test data
y_pred = rf_model.predict(X_test)
print("y_test shape:", y_test.shape)
y_train
```

หลังจากนั้นทำการประเมินความเที่ยงตรงด้วยการใช้ค่า Mean Squared Error และค่า R^2

```
# Evaluate the model (e.g., using mean squared error)
from sklearn.metrics import mean_squared_error

mse = mean_squared_error(y_test, y_pred)
print("Mean Squared Error:", mse)

# You can also evaluate using other metrics like R-squared, etc.
from sklearn.metrics import r2_score

r2 = r2_score(y_test, y_pred)
print("R-squared:", r2)
```

Mean Squared Error: 0.16466270309452602

R-squared: 0.839882499866233

จากค่า MSE และ R^2 โมเดลมีความเที่ยงตรงพอสมควร เพื่อสร้างความเข้าใจเพิ่มเติมในผลลัพธ์ของการพยากรณ์ ทำได้ด้วยการทำ visualization ระหว่างค่าพยากรณ์กับค่าจริงดังแสดงในตัวอย่างที่ 10.4 โดยจะทำการ ทดลอง plot กราฟเปรียบเทียบจากการสั่งซื้อ 50 ครั้งแรก โดยจะต้องทำการแปลงค่า log ยอดขายเป็นค่ายอดขายปกติก่อน

ตัวอย่าง 10.4 การทำ Visualization ค่าพยากรณ์กับค่าจริงของการสั่งซื้อ 50 ครั้งแรก

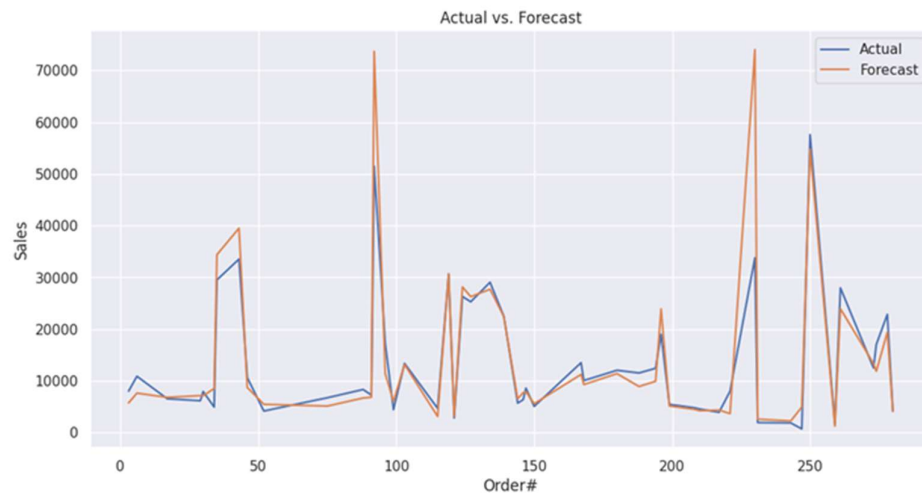
```
# add column to df_predict

df_predict = pd.DataFrame({'Actual': np.exp(y_test), 'Forecast': np.exp(y_pred)})

# sort df_predict in ascending order of index
df_predict = df_predict.sort_index(ascending=True)

# plot line graph of actual and forecast only 50 rows

plt.figure(figsize=(12, 6))
plt.plot(df_predict['Actual'].iloc[:50], label='Actual')
plt.plot(df_predict['Forecast'].iloc[:50], label='Forecast')
plt.xlabel('Index')
plt.ylabel('Value')
plt.title('Actual vs. Forecast')
plt.legend()
plt.show()
```



ภาพที่ 10.1 การเปรียบเทียบระหว่างค่าจริงกับค่าพยากรณ์

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

ภาพข้างต้นแสดงให้เห็นว่าค่าพยากรณ์มีค่าใกล้เคียงกับค่าจริง อย่างไรก็ตามเราสามารถปรับปรุง การพยากรณ์ให้มีค่าเที่ยงตรงมากยิ่งขึ้นโดยใช้ GridSearchCV ฟังก์ชัน ใน Scikit-learn Library การทำ GridSearchCV เพื่อหาค่า hyperparameters ของโมเดล ที่สามารถปรับปรุงความเที่ยงตรงของการพยากรณ์ให้ดีขึ้น แต่ผู้อ่านควรทำความเข้าใจว่าการทำ GridSearchCV จะใช้เวลานานกว่าที่คาดในการหาคำตอบเพราะต้องรันโปรแกรมสำหรับทุกๆ combinations ของ hyperparameters เพื่อหาค่าที่เหมาะสมที่สุด ขั้นตอนการทำ GridSearch แสดงดังต่อไปนี้

ตัวอย่าง 10.5 การทำ Visualization ค่าพยากรณ์กับค่าจริงของการสั่งซื้อ 50 ครั้งแรก

```
# do the gridsearchcv to improve the model

from sklearn.model_selection import GridSearchCV

# Define the parameter grid for GridSearchCV
param_grid = {
    'n_estimators': [50, 100, 200],
    'max_depth': [None, 10, 20],
    'min_samples_split': [2, 5, 10],
```

```
'min_samples_leaf': [1, 2, 4]
}

# Create a Random Forest Regressor model
rf_model = RandomForestRegressor(random_state=42)

# Create a GridSearchCV object
grid_search = GridSearchCV(estimator=rf_model, param_grid=param_grid,
                           scoring='neg_mean_squared_error', cv=5, n_jobs=-1)

# Fit the GridSearchCV object to the training data
grid_search.fit(X_train, y_train)

# Print the best parameters found by GridSearchCV
print("Best parameters:", grid_search.best_params_)

# Get the best model from GridSearchCV
best_rf_model = grid_search.best_estimator_

# Make predictions on the test data using the best model
y_pred = best_rf_model.predict(X_test)

# Evaluate the best model
mse = mean_squared_error(y_test, y_pred)
print("Mean Squared Error (Best Model):", mse)

r2 = r2_score(y_test, y_pred)
print("R-squared (Best Model):", r2)
```

Best parameters: {'max_depth': None, 'min_samples_leaf': 1, 'min_samples_split': 5, 'n_estimators': 200}

Mean Squared Error (Best Model): 0.16197311269291984

R-squared (Best Model): 0.8424978492039743

จากการทำ GridSearch พบว่าความเที่ยงตรงถูกปรับปรุงให้ดีขึ้น โดยค่า MSE. มีค่าลดลงจากเดิม 0.165 เป็น 0.162 และค่า R^2 เพิ่มขึ้นจาก 0.840 เป็น 0.842 และ model ที่ได้จากการทำ GridSearchCV ไปทำในขั้นตอนการวิเคราะห์ความสำคัญของฟีเจอร์ต่อไปได้ (feature importance) เพื่อศึกษาความสำคัญของฟีเจอร์ว่าฟีเจอร์ไหนมีความสำคัญต่อการพยากรณ์ยอดขายโดยสามารถทำได้ในตัวอย่าง 10.6

ตัวอย่าง 10.6 การทำ Feature Importance

```
# feature importance of model

# Get feature importances from the best model
feature_importances = best_rf_model.feature_importances_

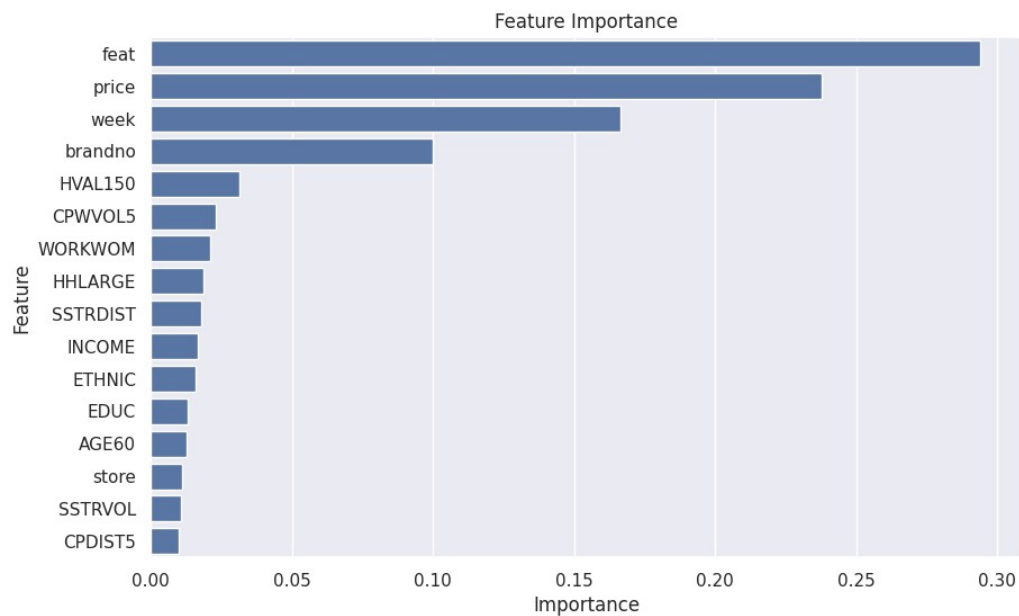
# Create a DataFrame to display feature importances
feature_importance_df = pd.DataFrame({'Feature': X_train.columns, 'Importance':
feature_importances})

# Sort the DataFrame by importance in descending order
feature_importance_df = feature_importance_df.sort_values('Importance', ascending=False)

# Print or display the feature importances
print(feature_importance_df)

# You can also visualize the feature importances using a bar chart
plt.figure(figsize=(10, 6))
sns.barplot(x='Importance', y='Feature', data=feature_importance_df)
plt.title('Feature Importance')
plt.xlabel('Importance')
plt.ylabel('Feature')
plt.show()
```

	Feature	Importance
3	feat	0.294014
4	price	0.237726
2	week	0.166376
1	brandno	0.099951
11	HVAL150	0.031392
15	CPWVOL5	0.022851
10	WORKWOM	0.021203
9	HHLARGE	0.018667
12	SSTRDIST	0.017896
8	INCOME	0.016719
7	ETHNIC	0.015939
6	EDUC	0.013030
5	AGE60	0.012699
0	store	0.011003
13	SSTRVOL	0.010512
14	CPDIST5	0.010023



ภาพที่ 10.2 การทำ Feature Importance เพื่อจัดลำดับความสำคัญของ Feature

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

จากการทำ Feature Importance พบว่า โมเดล Random Forest Regressor ให้ความสำคัญกับฟีเจอร์ที่สำคัญมากที่สุด 4 อันดับแรกได้แก่ feat (การโฆษณา) price (ราคา) week (สัปดาห์ที่) brandno (ยี่ห้อของน้ำส้ม)

1.4.2 XGboost

ในส่วนนี้เราจะมาลองทำการพยากรณ์ด้วยเทคนิคอื่นที่เรียกว่า XGboost ซึ่งมีหลักการพัฒนามาจากแนวคิดของ Gradient Boosting Machine (GBM) ซึ่งทำงานโดยการรวมเอาผลลัพธ์จากโมเดลต้นไม้ตัดสินใจ (decision trees) หลายๆ ต้นมาทำงานร่วมกันเพื่อเพิ่มความแม่นยำของการทำนาย โดย XGBoost มีคุณสมบัติที่โดดเด่นกว่า GBM แบบเดิมในหลายด้าน โดยมีกระบวนการทำงานดังขั้นตอนอธิบายดังต่อไปนี้

1. **การบูสต์ (Boosting)** เป็นเทคนิคการเรียนรู้แบบ iterative learning ที่สร้างโมเดลชุดหนึ่งขึ้นมาเป็นลำดับ แต่ละโมเดลจะพยายามแก้ไขข้อผิดพลาดที่โมเดลก่อนหน้านี้ทำได้ โมเดลใหม่ในแต่ละรอบจะให้ความสำคัญกับตัวอย่างข้อมูลที่โมเดลก่อนหน้านี้ทำนายผิดพลาด นอกจากนั้นยังใช้แนวคิดของการลดค่า loss function (เช่น Mean Squared Error หรือ Log Loss) โดยการสร้างต้นไม้ตัดสินใจเพิ่มเติมที่ช่วยลดความคลาดเคลื่อนของผลลัพธ์
2. **Regularization** XGBoost เพิ่มการควบคุมความซับซ้อนของโมเดล (regularization) เพื่อป้องกันการ overfitting โดยมีการใช้ทั้ง L1 และ L2 regularization เพื่อควบคุมน้ำหนักของฟีเจอร์ในโมเดล
3. **Tree Pruning** ใช้วิธี "max_depth" ในการตัดต้นไม้เพื่อควบคุมโครงสร้างของต้นไม้ ลดความซับซ้อน และเพิ่มความเร็วในการคำนวณ
4. **การคำนวณแบบขนาน (Parallel Computing)** XGBoost ออกแบบมาให้รองรับการทำงานแบบขนาน จึงสามารถสร้างต้นไม้ได้อย่างรวดเร็วเมื่อเทียบกับ GBM แบบเดิม
5. **การจัดการ Missing Values** XGBoost สามารถจัดการกับข้อมูลที่ขาดหายได้อย่างมีประสิทธิภาพ โดยจะพิจารณาเส้นทางที่เหมาะสมที่สุดสำหรับค่าที่ขาดหายระหว่างการสร้างต้นไม้

ตัวอย่างที่ 10.7 แสดงขั้นตอนการใช้เทคนิค XGboost พยากรณ์อุปสงค์ของชุดข้อมูลกรณีศึกษาเดียวกัน ด้วย PYTHON อย่างละเอียดและแสดงโค้ดประกอบการอธิบายดังต่อไปนี้

ตัวอย่าง 10.7 การใช้วิธี XGboost

การแบ่งชุดข้อมูล training / testing และการสร้าง DMatrix สำหรับใช้วิเคราะห์ข้อมูลซึ่งจะช่วยให้ประมวลผลได้รวดเร็วกว่าโครงสร้างข้อมูลแบบอื่น นอกจากนั้น DMatrix ยังรองรับการทำงานแบบกระจาย (distributed training) และแบบขนาน (parallel processing) ทำให้เหมาะสมกับเทคนิค XGboost โดยในขั้นตอนนี้สามารถโค้ดด้วย PYTHON แสดงดังต่อไปนี้

```
# Target variable is 'logmove'
# and we use all other columns as features
X = data.drop(['brand', 'logmove'], axis=1)
y = data['logmove']

# Split the data into training and testing sets
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)

# Create DMatrix for XGBoost
dtrain = xgb.DMatrix(X_train, label=y_train)
dtest = xgb.DMatrix(X_test, label=y_test)
evallist = [(dtest, 'eval'), (dtrain, 'train')]
```

หลังจากนั้นจะกำหนดพารามิเตอร์ที่จำเป็นได้แก่

- a. **objective** กำหนดเป้าหมายของการเรียนรู้ เช่น
 - 1) reg:squarederror (สำหรับการถดถอย)
 - 2) binary:logistic (สำหรับการจัดประเภทแบบ binary)
 - 3) multi:softprob (สำหรับการจัดประเภทหลายคลาส)
- b. **eval_metric** ตัววัดผลสำหรับการประเมิน เช่น rmse, mae, logloss, error, auc
- c. **seed** ค่า seed สำหรับการสุ่ม สามารถใช้ค่า [0, ∞]

```
# Define XGBoost parameters
params = {"booster": "gbtree",
          'objective': 'reg:squarederror', # Regression task
          'eval_metric': 'rmse', # Root Mean Squared Error
          'eta': 0.03, # Learning rate
```

```
'max_depth': 8, # Maximum tree depth
'subsample': 0.95, # Subsample ratio of the training instance
'colsample_bytree': 0.8, # Subsample ratio of columns when constructing each tree
}
```

ตัวอย่าง 10.8 การ Train โมเดลการพยากรณ์และประเมินความเที่ยงตรง

```
# Train XGBoost model
num_rounds = 6000 # Number of boosting rounds
model = xgb.train(params, dtrain, num_rounds, evals=evallist, early_stopping_rounds=100)

# Make predictions on the test set
y_pred = model.predict(dtest)

# Import necessary libraries
from sklearn.metrics import mean_squared_error, r2_score # Import mean_squared_error and
r2_score

# Evaluate the model
mse = mean_squared_error(y_test, y_pred)
print("Mean Squared Error:", mse)

r2 = r2_score(y_test, y_pred)
print("R-squared:", r2)

rmspe = rmse_percentage_error(y_test, y_pred)
print("Root Mean Square Percentage Error:", rmspe)
```

R-squared: 0.8827166041350083

Root Mean Square Percentage Error: 0.0408496831196627

จากผลลัพธ์ของการประเมินโมเดลที่ใช้เทคนิค XGboost พบว่า มีค่า MSE. เท่ากับ 0.121 และค่า R^2 0.883 และค่า RMSPE. เท่ากับ 0.041 แสดงให้เห็นว่ามีความเที่ยงตรงมากกว่าการใช้เทคนิค Random Forest

10.5 สรุปผลลัพธ์จากการพยากรณ์อุปสงค์

จากการพยากรณ์อุปสงค์โดยใช้เทคนิคการเรียนรู้ของเครื่อง ได้แก่ Random Forest และ XGBoost พบว่าแบบจำลองทั้งสองสามารถพยากรณ์ยอดขายได้อย่างมีประสิทธิภาพ โดย Random Forest ให้ผลลัพธ์ในระดับที่น่าพอใจ โดยมีค่า Mean Squared Error (MSE) ประมาณ 0.165 และค่า R-squared (R^2) ประมาณ 0.840 และเมื่อปรับแต่งพารามิเตอร์ด้วย grid search พบว่าประสิทธิภาพของโมเดลดีขึ้น โดยค่า MSE ลดลงเหลือประมาณ 0.162 และค่า R^2 เพิ่มขึ้นเป็นประมาณ 0.842 อีกทั้งการวิเคราะห์ความสำคัญของตัวแปร (feature importance) ยังแสดงให้เห็นว่าปัจจัยสำคัญที่มีอิทธิพลต่อยอดขาย ได้แก่ การโฆษณา (FEAT) ราคา (PRICE) ช่วงเวลา (WEEK) และยี่ห้อสินค้า (BRAND) ในขณะที่แบบจำลอง XGBoost ให้ผลลัพธ์ที่มีความแม่นยำสูงกว่า โดยมีค่า MSE ประมาณ 0.121 ค่า R^2 ประมาณ 0.883 และค่า RMSPE ประมาณ 0.041 ซึ่งแสดงให้เห็นถึงความสามารถของโมเดลในการลดความคลาดเคลื่อนและจับรูปแบบข้อมูลที่ซับซ้อนได้ดีกว่า Random Forest

อย่างไรก็ตาม เมื่อเปรียบเทียบกับวิธีการพยากรณ์แบบดั้งเดิม เช่น Moving Average Exponential Smoothing และ ARIMA พบว่าวิธีดั้งเดิมมีข้อเสียเปรียบหลายประการ โดยเฉพาะข้อจำกัดในการจัดการกับข้อมูลที่มีความซับซ้อนและไม่เป็นเชิงเส้น เนื่องจากแบบจำลองส่วนใหญ่ตั้งอยู่บนสมมติฐานเชิงเส้น (linear assumption) จึงไม่สามารถสะท้อนความสัมพันธ์ของตัวแปรหลายมิติ เช่น ราคา โปรโมชัน และปัจจัยด้านประชากรศาสตร์ได้อย่างครบถ้วน นอกจากนี้ วิธีดั้งเดิมมักรองรับข้อมูลแบบตัวแปรเดียว (univariate) เป็นหลัก ทำให้ไม่สามารถนำข้อมูลจากหลายแหล่งมาประยุกต์ใช้ร่วมกันได้อย่างมีประสิทธิภาพ อีกทั้งยังมีข้อจำกัดในการปรับตัวต่อการเปลี่ยนแปลงของตลาดที่รวดเร็ว เช่น พฤติกรรมผู้บริโภคที่เปลี่ยนแปลงหรือการแข่งขันที่รุนแรง

โดยสรุป ผลการศึกษาชี้ให้เห็นว่าเทคนิคการเรียนรู้ของเครื่อง โดยเฉพาะ XGBoost มีประสิทธิภาพเหนือกว่าวิธีดั้งเดิมในการพยากรณ์อุปสงค์ในกรณีที่ข้อมูลมีความซับซ้อนและมีหลายปัจจัยร่วมกัน อย่างไรก็ตาม การเลือกใช้วิธีการพยากรณ์ยังควรพิจารณาจากลักษณะของข้อมูลและบริบทของธุรกิจ เพื่อให้ได้ผลลัพธ์ที่เหมาะสมที่สุดในการสนับสนุนการตัดสินใจในซัพพลายเชน

เอกสารอ้างอิง

- Areerakulkan, N., Moryadee, C., Trakoonsanti, L., Khaengkhan, M., & Setthachotsombut, N. (2024). *A new product's demand forecasting using artificial neural network*. In *Proceedings of the 26th International Conference on Enterprise Information Systems (ICEIS 2024)* (Vol. 1, pp. 417-424). SciTePress. <https://doi.org/10.5220/0012734800003690>.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5-32. <https://doi.org/10.1023/A:1010933404324>.
- Carbonneau, R., Laframboise, K., & Vahidov, R. (2008). Application of machine learning techniques for supply chain demand forecasting. *European Journal of Operational Research*, 184(3), 1140-1154. <https://doi.org/10.1016/j.ejor.2006.12.004>.
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785-794). <https://doi.org/10.1145/2.939672.2939785>
- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2018). The M4 competition: Results, findings, conclusion and way forward. *International Journal of Forecasting*, 34(4), 802-808. <https://doi.org/10.1016/j.ijforecast.2018.06.001>
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77-84. <https://doi.org/10.1111/jbl.12010>

บทที่ 11

จากทฤษฎีสู่การปฏิบัติ: การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่อง เพื่อการจัดเส้นทางขนส่ง

From Theory to Practice: Data Analytics and Machine Learning for Transportation Routing

หัวข้อ

- 11.1 วัตถุประสงค์ของกรณีศึกษา
- 11.2 บริบทธุรกิจ
- 11.3 ชุดข้อมูลและตัวแปรที่ใช้
- 11.4 ขั้นตอนการวิเคราะห์ด้วยการเรียนรู้ของเครื่อง
- 11.5 สรุปผลลัพธ์จากการจัดเส้นทางขนส่ง

การจัดเส้นทางขนส่ง (transportation routing) เป็นหนึ่งในปัญหาสำคัญในงานด้านโลจิสติกส์และการจัดการซัพพลายเชน โดยมีเป้าหมายเพื่อกำหนดเส้นทางขนส่งที่มีประสิทธิภาพสูงสุดภายใต้ข้อจำกัดต่าง ๆ เช่น เวลา ต้นทุน ความจุของยานพาหนะ และความต้องการของลูกค้า ปัญหานี้มักถูกจัดอยู่ในกลุ่มของปัญหา vehicle routing problem (VRP) ซึ่งเป็นปัญหาที่มีความซับซ้อนสูงและจัดอยู่ในประเภท NP-hard ทำให้การหาคำตอบที่เหมาะสมที่สุด (optimal solution) ด้วยวิธีดั้งเดิมทำได้ยาก โดยเฉพาะเมื่อขนาดของปัญหาเพิ่มขึ้น (Toth & Vigo, 2014)

ในอดีตการแก้ปัญหาการจัดเส้นทางมักอาศัยวิธีการเชิงคณิตศาสตร์และอัลกอริทึมฮิวริสติก เช่น Clarke-Wright Savings Algorithm หรือ Genetic Algorithm ซึ่งสามารถให้คำตอบที่ดีในเวลาอันจำกัด แต่ยังมีข้อจำกัดในการจัดการกับข้อมูลขนาดใหญ่และสภาพแวดล้อมที่มีความไม่แน่นอน (Laporte, 2009) อย่างไรก็ตาม การพัฒนาเทคโนโลยีด้านข้อมูลและพลังการประมวลผลในปัจจุบันทำให้แนวคิดของการเรียนรู้ของเครื่อง (ML) เข้ามามีบทบาทสำคัญในการแก้ปัญหานี้

ML ช่วยให้ระบบสามารถเรียนรู้รูปแบบจากข้อมูลในอดีต เช่น รูปแบบการจราจร พฤติกรรมการสั่งซื้อ หรือเวลาการส่งมอบ เพื่อปรับปรุงการตัดสินใจในการจัดเส้นทางให้มีประสิทธิภาพมากยิ่งขึ้น โดยเฉพาะเทคนิคขั้นสูง เช่น Deep Learning และ Reinforcement Learning ที่สามารถใช้

ในการเรียนรู้เชิงลำดับและการตัดสินใจแบบไดนามิก ซึ่งเหมาะสมกับปัญหาการจัดเส้นทางที่มีการเปลี่ยนแปลงตลอดเวลา (Kool et al., 2019; Nazari et al., 2018)

นอกจากนี้ อีกหนึ่งเทคนิคที่ได้รับความสนใจในการประยุกต์ใช้กับปัญหาการจัดเส้นทางคือ Spectral Clustering (SC) ซึ่งเป็นวิธีการจัดกลุ่มข้อมูลที่อาศัยโครงสร้างของกราฟและค่าลักษณะเฉพาะ (Eigenvalues) ของเมทริกซ์ Laplacian เพื่อแบ่งกลุ่มจุดข้อมูลที่มีความคล้ายคลึงกัน เทคนิคนี้เหมาะอย่างยิ่งสำหรับการจัดกลุ่มลูกค้าหรือจุดส่งสินค้าในเชิงพื้นที่ก่อนนำไปวางแผนเส้นทาง (Cluster-first, Route-second approach) โดยสามารถลดความซับซ้อนของปัญหา vehicle routing ได้อย่างมีประสิทธิภาพ (Von Luxburg, 2007) การใช้ SC ในบริบทของการขนส่งช่วยให้สามารถจับโครงสร้างเชิงพื้นที่ที่ซับซ้อน เช่น กลุ่มลูกค้าที่มีการกระจายตัวไม่เป็นรูปทรงเรขาคณิตแบบง่าย ซึ่งวิธีการจัดกลุ่มแบบดั้งเดิม เช่น K-means อาจไม่สามารถจัดการได้ดี ส่งผลให้การแบ่งโซนการให้บริการ (service zones) มีความแม่นยำมากขึ้น และนำไปสู่การลดระยะทางรวมของการขนส่งและเวลาในการจัดส่ง (Shi & Malik, 2000; Von Luxburg, 2007) เมื่อผสาน SC เข้ากับเทคนิค ML และแบบจำลอง Optimization จะช่วยเพิ่มศักยภาพในการจัดการปัญหาการขนส่งที่มีความซับซ้อนสูง โดยเฉพาะในสภาพแวดล้อมที่มีข้อมูลจำนวนมากและมีความไม่แน่นอน เช่น เมืองขนาดใหญ่หรือระบบขนส่งอัจฉริยะ ซึ่งสามารถนำไปสู่การลดต้นทุน เพิ่มประสิทธิภาพ และยกระดับคุณภาพการให้บริการได้อย่างยั่งยืน (Waller & Fawcett, 2013)

ดังนั้น กรณีสึกษาการจัดเส้นทางขนส่งด้วยวิธีการเรียนรู้ของเครื่องจึงเป็นแนวทางที่มีศักยภาพสูงในการพัฒนาระบบโลจิสติกส์สมัยใหม่ โดยอาศัยการผสมผสานระหว่างการเรียนรู้จากข้อมูลและการเพิ่มประสิทธิภาพเชิงคณิตศาสตร์ เพื่อรองรับความท้าทายของซัพพลายเชนในยุคดิจิทัล

11.1 วัตถุประสงค์ของกรณีสึกษา

กรณีสึกษานี้อ้างอิงจากการศึกษาของ Areerakulkan et al. (2026) มีวัตถุประสงค์เพื่อศึกษาการประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่องร่วมกับวิธีการเพิ่มประสิทธิภาพในการแก้ปัญหาการจัดเส้นทางขนส่ง ภายใต้กรอบของ VRP ซึ่งเป็นปัญหาที่มีความซับซ้อนและมีข้อจำกัดหลายด้าน โดยเฉพาะข้อจำกัดด้านความจุของยานพาหนะ ในกรณีสึกษานี้ได้มุ่งเน้นการนำ SC มาใช้ในการจัดกลุ่มลูกค้าตามความใกล้เคียงเชิงพื้นที่ เพื่อช่วยลดความซับซ้อนของปัญหา ก่อนนำไปวางแผนเส้นทางขนส่งในแต่ละกลุ่ม นอกจากนี้ กรณีสึกษายังมีเป้าหมายเพื่อให้เข้าใจกระบวนการตั้งแต่การเตรียมข้อมูลตำแหน่งทางภูมิศาสตร์ของจุดส่งสินค้าและคลังสินค้า การคำนวณระยะทางระหว่างจุดต่าง ๆ และการสร้าง distance matrix ซึ่งเป็นพื้นฐานสำคัญในการวางแผนเส้นทาง จากนั้นจึงนำวิธีการจัดเส้นทางแบบง่าย เช่น Nearest Neighbor มาประยุกต์ใช้ร่วมกับผลลัพธ์จากการจัดกลุ่ม เพื่อสร้างเส้นทางขนส่งที่เหมาะสมภายใต้ข้อจำกัดที่กำหนด ทำให้เพิ่มประสิทธิภาพในการให้บริการ และการ

สนับสนุนการตัดสินใจเชิงข้อมูล ซึ่งสามารถนำไปประยุกต์ใช้ในบริบทของซัพพลายเชนสมัยใหม่ได้อย่างมีประสิทธิภาพ

11.2 บริบทธุรกิจ

ในยุคที่การแข่งขันทางธุรกิจมีความเข้มข้นมากขึ้น องค์กรด้านโลจิสติกส์และการกระจายสินค้าจำเป็นต้องพัฒนาประสิทธิภาพในการดำเนินงานอย่างต่อเนื่อง โดยเฉพาะในส่วนของการจัดส่งสินค้า (last-mile delivery) ซึ่งเป็นขั้นตอนที่มีต้นทุนสูงและมีผลกระทบโดยตรงต่อความพึงพอใจของลูกค้า กรณีศึกษานี้ตั้งอยู่บนบริบทของบริษัทผู้ให้บริการขนส่งสินค้าภายในเมือง ที่ต้องจัดส่งสินค้าไปยังลูกค้าหลายจุดในแต่ละวัน ภายใต้ข้อจำกัดด้านจำนวนยานพาหนะ ความจุของรถ และระยะเวลาการจัดส่ง

บริษัทมีคลังสินค้า (depot) เป็นจุดเริ่มต้นของการขนส่ง และมีลูกค้ากระจายตัวอยู่ในพื้นที่เมืองในลักษณะที่ไม่เป็นระเบียบ ทำให้การวางแผนเส้นทางมีความซับซ้อน หากใช้วิธีการวางแผนแบบดั้งเดิม อาจทำให้เกิดปัญหา เช่น ระยะทางรวมสูงเกินไป การใช้ทรัพยากรไม่คุ้มค่า หรือการส่งสินค้าล่าช้า ซึ่งส่งผลกระทบต่อทั้งต้นทุนและระดับการให้บริการ เพื่อแก้ไขปัญหาดังกล่าว บริษัทจึงนำแนวคิดของ VRP มาใช้เป็นกรอบในการวิเคราะห์ พร้อมทั้งประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่อง โดยเฉพาะ SC ในการจัดกลุ่มลูกค้าตามความใกล้เคียงเชิงพื้นที่ เพื่อลดความซับซ้อนของการวางแผนเส้นทาง และช่วยให้สามารถจัดสรรงานให้กับยานพาหนะแต่ละคันได้อย่างเหมาะสม

ภายหลังจากการจัดกลุ่มลูกค้า บริษัทใช้วิธีการสร้างเส้นทางขนส่งภายในแต่ละกลุ่ม โดยคำนึงถึงข้อจำกัดด้านความจุของรถและระยะทาง ซึ่งบริษัทมีรถบรรทุกขนส่งภายในกรุงเทพฯ อยู่ 4 คัน โดยแต่ละคันจะถูกกำหนดให้มีความจุมากที่สุด 120 หน่วย เพื่อให้ได้เส้นทางที่มีประสิทธิภาพมากขึ้น แนวทางนี้ช่วยให้บริษัทสามารถลดระยะทางรวมในการขนส่ง ลดต้นทุนเชื้อเพลิง และเพิ่มความตรงต่อเวลาในการจัดส่งสินค้า ดังนั้นกรณีศึกษานี้สะท้อนให้เห็นถึงความท้าทายจริงในงานโลจิสติกส์สมัยใหม่ และแสดงให้เห็นถึงบทบาทสำคัญของการผสมผสานเทคโนโลยี ML กับการเพิ่มประสิทธิภาพเชิงคณิตศาสตร์ในการยกระดับขีดความสามารถในการแข่งขันขององค์กรในซัพพลายเชน

11.3 ชุดข้อมูลและตัวแปรที่ใช้

ชุดข้อมูลที่ใช้ในกรณีศึกษานี้เป็นข้อมูลเชิงพื้นที่ (geospatial data) ที่ถูกจัดเก็บในรูปแบบตาราง โดยแต่ละแถวแทนจุดให้บริการหนึ่งจุด ซึ่งประกอบด้วยทั้งคลังสินค้า (depot) และลูกค้า (customers) รวมทั้งหมด 51 จุด โดยข้อมูลทั้งหมดตั้งอยู่ในพื้นที่ กรุงเทพมหานคร และกระจายตัวอยู่ในหลายเขตสำคัญ เช่น ดุสิต ดินแดง ธนบุรี ตลิ่งชัน ภาษีเจริญ ปทุมวัน คลองเตย บางบอน และบางกะปิ ข้อมูลดังกล่าวทำหน้าที่เป็นอินพุตหลักในการวิเคราะห์และแก้ปัญหาการจัดเส้นทาง

ขนส่งภายใต้กรอบของ VRP โดยมีการออกแบบตัวแปรให้ครอบคลุมทั้งมิติด้านตำแหน่ง ปริมาณความต้องการ และตัวแปรที่ได้จากการประมวลผลเพิ่มเติม

ในส่วนของตัวแปรเชิงตำแหน่ง ข้อมูลถูกกำหนดในรูปแบบพิกัดละติจูด (latitude) และลองจิจูด (longitude) ของแต่ละจุด ซึ่งใช้สำหรับระบุตำแหน่งทางภูมิศาสตร์ของคลังสินค้าและลูกค้า ในเขตกรุงเทพมหานคร และเป็นพื้นฐานสำคัญในการคำนวณระยะทางระหว่างจุดต่าง ๆ รวมถึงการจัดกลุ่มลูกค้าตามความใกล้เคียงเชิงพื้นที่ โดยในกรณีศึกษาได้ประยุกต์ใช้ SC เพื่อแบ่งกลุ่มลูกค้าออกเป็นหลายกลุ่มให้เหมาะสมกับการวางแผนเส้นทาง นอกจากนี้ ยังมีตัวแปรด้านความต้องการ ซึ่งแสดงถึงปริมาณสินค้าที่ลูกค้าแต่ละรายต้องการ ตัวแปรนี้มีบทบาทสำคัญในการกำหนดข้อจำกัดของปัญหา เนื่องจากต้องควบคุมให้ผลรวมของความต้องการในแต่ละเส้นทางไม่เกินความจุของยานพาหนะที่กำหนดไว้

จากข้อมูลพื้นฐานดังกล่าว ได้มีการสร้างตัวแปรเพิ่มเติมเพื่อใช้ในการวิเคราะห์ ได้แก่ เมตริกระยะทาง ซึ่งคำนวณจากพิกัดด้วยวิธี Haversine Distance Function เพื่อใช้เป็นข้อมูลนำเข้าในการจัดเส้นทางและการคำนวณหาระยะทางรวมของแต่ละเส้นทาง ซึ่งจะถูกใช้ในกระบวนการจัดกลุ่มแบบ SC รวมถึงตัวแปรผลลัพธ์ เช่น cluster label ที่ระบุว่าลูกค้าแต่ละรายอยู่ในกลุ่มใด และ route sequence ซึ่งแสดงลำดับการเดินทางของยานพาหนะในแต่ละเส้นทาง

โดยภาพรวม ชุดข้อมูลและตัวแปรที่ใช้ในกรณีศึกษานี้สะท้อนบริบทของการขนส่งในเขตเมืองของกรุงเทพมหานคร ซึ่งมีความหนาแน่นและความซับซ้อนสูง และเป็นตัวอย่างที่เหมาะสมสำหรับการประยุกต์ใช้เทคนิค ML ร่วมกับการเพิ่มประสิทธิภาพในการวางแผนเส้นทางขนส่งในงานโลจิสติกส์จริง

11.4 ขั้นตอนการวิเคราะห์ด้วยวิธีผสมระหว่าง Spectral Clustering และ Nearest Neighbor Algorithms

ขั้นตอนการวิเคราะห์เริ่มจากการ import packages และข้อมูลที่ต้องใช้ดังต่อไปนี้

ตัวอย่าง 11.1 การ import packages ที่ต้องใช้และข้อมูลนำเข้าเพื่อเริ่มทำการวิเคราะห์

การ import package และข้อมูลที่จำเป็นต้องใช้พร้อมทั้งแสดง 5 แถวแรก

```
# conduct vrp problem with capacity constraint using spectral clustering and nearest neighbor
import pandas as pd
```

```
import numpy as np
import matplotlib.pyplot as plt
from sklearn.cluster import SpectralBiclustering, SpectralClustering
from scipy.spatial.distance import cdist

import pandas as pd
# read location.csv file

df = pd.read_csv('location.csv')
# show the first five rows

df.head()
```

	Name	District	Latitude	Longitude	Demand
0	DP1	Din Daeng	13.778	100.567	9
1	Depot	Dusit	13.758	100.517	13
2	DP2	Thon Buri	13.731	100.490	9
3	DP3	Thon Buri	13.737	100.493	6
4	DP4	Taling Chan	13.760	100.456	9

หลังจากนั้นทำการ plot แสดงตำแหน่งของจุดส่งของและDepot

ตัวอย่าง 11.2 การ plot ตำแหน่งจุดส่งของและDepot

ทำการ Plot แสดงตำแหน่งของ Depot (รูปดาว) และจุดส่งของต่างๆตามพิกัด ละติจูด ลองจิจูด

```
import matplotlib.pyplot as plt

# PLOT scatter diagram and only for Depot will have star label

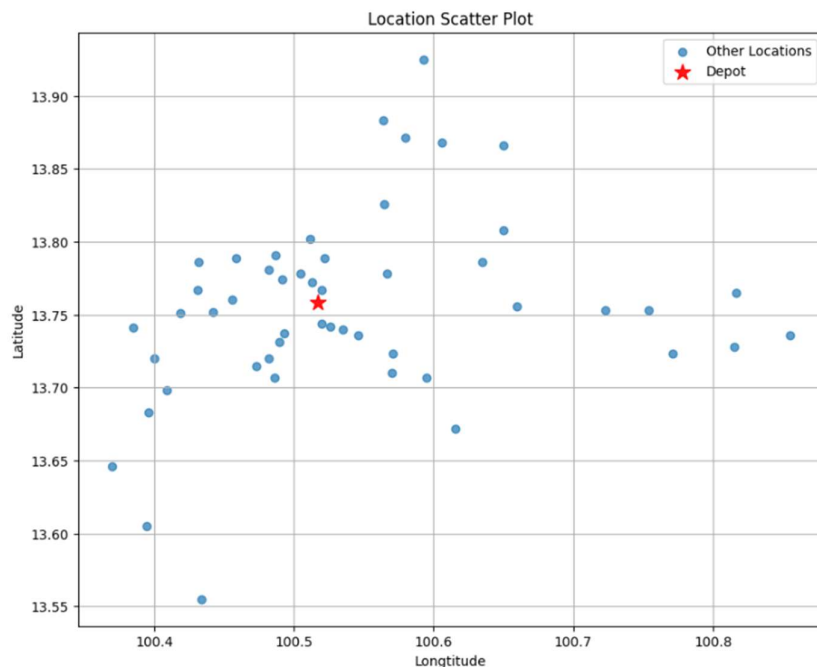
# Assuming 'df' has columns 'latitude', 'longitude', and 'LocationType'
# where 'LocationType' can be 'Depot' or other types

plt.figure(figsize=(10, 8))
```

```
# Plot other locations with default markers
other_locations = df[df['Name'] != 'Depot']
plt.scatter(other_locations['Longitude'], other_locations['Latitude'], label='Other Locations',
            alpha=0.7)

# Plot Depot locations with star markers
depot_locations = df[df['Name'] == 'Depot']
plt.scatter(depot_locations['Longitude'], depot_locations['Latitude'], marker='*', s=150,
            color='red', label='Depot', alpha=0.9) # s is marker size

plt.xlabel('Longitude')
plt.ylabel('Latitude')
plt.title('Location Scatter Plot')
plt.legend()
plt.grid(True)
plt.show()
```



ภาพที่ 11.1 ตำแหน่งของ Depot (★) กับ จุดส่งของ (●)

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

ตัวอย่าง 11.3 ทำการจัดกลุ่มจุดส่งของด้วยวิธี Spectral Clustering

เราสามารถทำการจัดกลุ่มด้วยวิธี Spectral Clustering ดังต่อไปนี้

```
model = SpectralClustering(n_clusters=4, affinity='nearest_neighbors', random_state=0)
df['Cluster'] = model.fit_predict(locations)
MAX_CAPACITY = 120
print("Spectral clustering performed and MAX_CAPACITY set.")
print(df.head())
```

First 5 rows of df_reassigned with updated clusters:

	Name	District	Latitude	Longitude	Demand	Cluster
0	DP1	Din Daeng	13.778	100.567	9	0
1	Depot	Dusit	13.758	100.517	13	3
2	DP2	Thon Buri	13.731	100.490	9	3
3	DP3	Thon Buri	13.737	100.493	6	3
4	DP4	Taling Chan	13.760	100.456	9	1

ตัวอย่าง 11.4 ทำการ Visualize Updated Clusters

ทำการ Visualize Updated Clusters ดังต่อไปนี้

```
import matplotlib.pyplot as plt
import seaborn as sns

plt.figure(figsize=(12, 10))

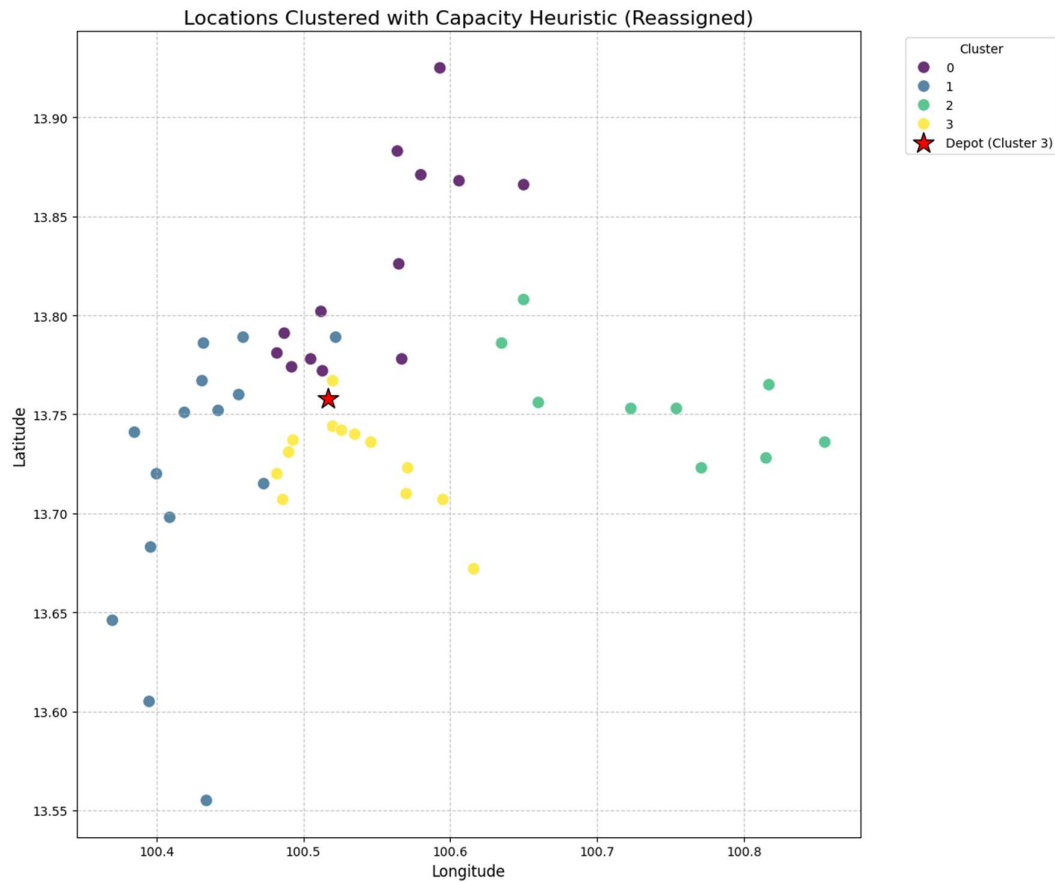
# Plot regular points, colored by cluster
sns.scatterplot(
    data=df_reassigned[df_reassigned['Name'] != 'Depot'],
    x='Longitude',
    y='Latitude',
    hue='Cluster',
    palette='viridis', # Or any other suitable palette
    s=100, # Size of regular points
```

```
alpha=0.8,
legend='full'
)

# Highlight the Depot location
depot_location = df_reassigned[df_reassigned['Name'] == 'Depot']
plt.scatter(
    depot_location['Longitude'],
    depot_location['Latitude'],
    marker='*', # Star marker
    s=300, # Larger size for emphasis
    color='red', # Red color for Depot
    edgecolor='black', # Black edge for better visibility
    label='Depot (Cluster ' + str(depot_location['Cluster'].iloc[0]) + ')',
    zorder=5 # Ensure depot is on top
)

plt.title('Locations Clustered with Capacity Heuristic (Reassigned)', fontsize=16)
plt.xlabel('Longitude', fontsize=12)
plt.ylabel('Latitude', fontsize=12)
plt.legend(title='Cluster', bbox_to_anchor=(1.05, 1), loc='upper left')
plt.grid(True, linestyle='--', alpha=0.7)
plt.tight_layout()
plt.show()
```

บทที่ 11 การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อการจัดเส้นทางขนส่ง



ภาพที่ 11.2 การจัดกลุ่มด้วย Spectral Clustering แบ่งตามสีทั้ง 4 สี

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

ตัวอย่าง 11.5 ทำการจัดเส้นทางในแต่ละ Cluster ด้วยวิธี Nearest Neighbor

เราสามารถทำการจัดเส้นทางในแต่ละ cluster ด้วยวิธี Nearest Neighbor หลังจากนั้นจึงทำการ visualize เส้นทางได้ดังต่อไปนี้

```
import matplotlib.pyplot as plt
import seaborn as sns
# Removed from scipy.spatial.distance import euclidean as it's no longer needed
import numpy as np
```

```
# 1. Initialize dictionaries
cluster_routes = {}
total_route_distances = {}

# Iterate through each cluster
for cluster_id, customer_points_df in cluster_customer_points.items():
    print(f"\nProcessing Cluster {cluster_id}...")

    # 2a. Define the starting point of the route as the 'Depot'
    depot_point = {
        'Name': 'Depot',
        'Latitude': depot_latitude,
        'Longitude': depot_longitude,
        'Demand': df_reassigned[df_reassigned['Name'] == 'Depot']['Demand'].iloc[0]
    }

    # 2b. Create a current_route list, starting with the Depot's information
    current_route = [depot_point]

    # 2c. Create a working copy of the customer points for the current cluster
    unvisited_points_df = customer_points_df.copy()

    # 2d. Initialize current_total_distance to 0
    current_total_distance = 0.0

    # Set the current_point to the Depot for the first leg
    current_point = depot_point

    # 2e. While there are still unvisited points
    while not unvisited_points_df.empty:
        # Using haversine_distance instead of euclidean
        distances = unvisited_points_df.apply(
            lambda row: haversine_distance(
                current_point['Latitude'], current_point['Longitude'],
                row['Latitude'], row['Longitude']
            ),
```

```

        axis=1
    )

    # 2e.ii. Identify the next_point (the unvisited point with the minimum distance)
    next_point_idx = distances.idxmin()
    next_point = unvisited_points_df.loc[next_point_idx].to_dict()

    # 2e.iii. Add the next_point's information to the current_route list
    current_route.append(next_point)

    # 2e.iv. Add the calculated distance to current_total_distance
    current_total_distance += distances.min()

    # 2e.v. Remove the next_point from unvisited_points_df
    unvisited_points_df = unvisited_points_df.drop(next_point_idx)

    # 2e.vi. Update the current_point to the next_point
    current_point = next_point

    # 2f. After all customer points are visited, calculate the distance from the current_point
    back to the Depot
    # Using haversine_distance instead of euclidean
    distance_to_depot = haversine_distance(
        current_point['Latitude'], current_point['Longitude'],
        depot_latitude, depot_longitude
    )
    current_total_distance += distance_to_depot

    # 2g. Add the Depot's information again to the current_route list to close the loop
    current_route.append(depot_point)

    # 2h. Store the current_route in the cluster_routes dictionary
    cluster_routes[cluster_id] = current_route

    # 2i. Store the current_total_distance in the total_route_distances dictionary
    total_route_distances[cluster_id] = current_total_distance

```

```
print("Nearest Neighbor routes generated for all clusters (using Haversine distance).")
print("\nTotal Route Distances (Haversine):")
for cluster_id, distance in total_route_distances.items():
    print(f"Cluster {cluster_id}: {distance:.4f} km")

# 3. Create a scatter plot using matplotlib.pyplot and seaborn
plt.figure(figsize=(14, 12))

# 3a. Plot all customer locations from df_reassigned (excluding the 'Depot'), coloring them by
their 'Cluster' assignment
sns.scatterplot(
    data=df_reassigned[df_reassigned['Name'] != 'Depot'],
    x='Longitude',
    y='Latitude',
    hue='Cluster',
    palette='tab10',
    s=100,
    alpha=0.8,
    legend='full',
    zorder=2
)

# 3b. Plot the 'Depot' location as a distinct marker
depot_location = df_reassigned[df_reassigned['Name'] == 'Depot']
plt.scatter(
    depot_location['Longitude'],
    depot_location['Latitude'],
    marker='s', # Square marker for depot
    s=300,
    color='black',
    edgecolor='white',
    label='Depot',
    zorder=3
)
```

```
# 3c. For each route stored in cluster_routes:
# Define a list of colors for routes to make them distinct, if 'tab10' palette doesn't have
enough distinct colors for all routes
route_colors = plt.colormaps["Dark2"](np.linspace(0, 1, len(cluster_routes))) # Updated to use
plt.colormaps

for i, (cluster_id, route) in enumerate(cluster_routes.items()):
    route_coords = np.array([[point["Longitude"], point["Latitude"]] for point in route])

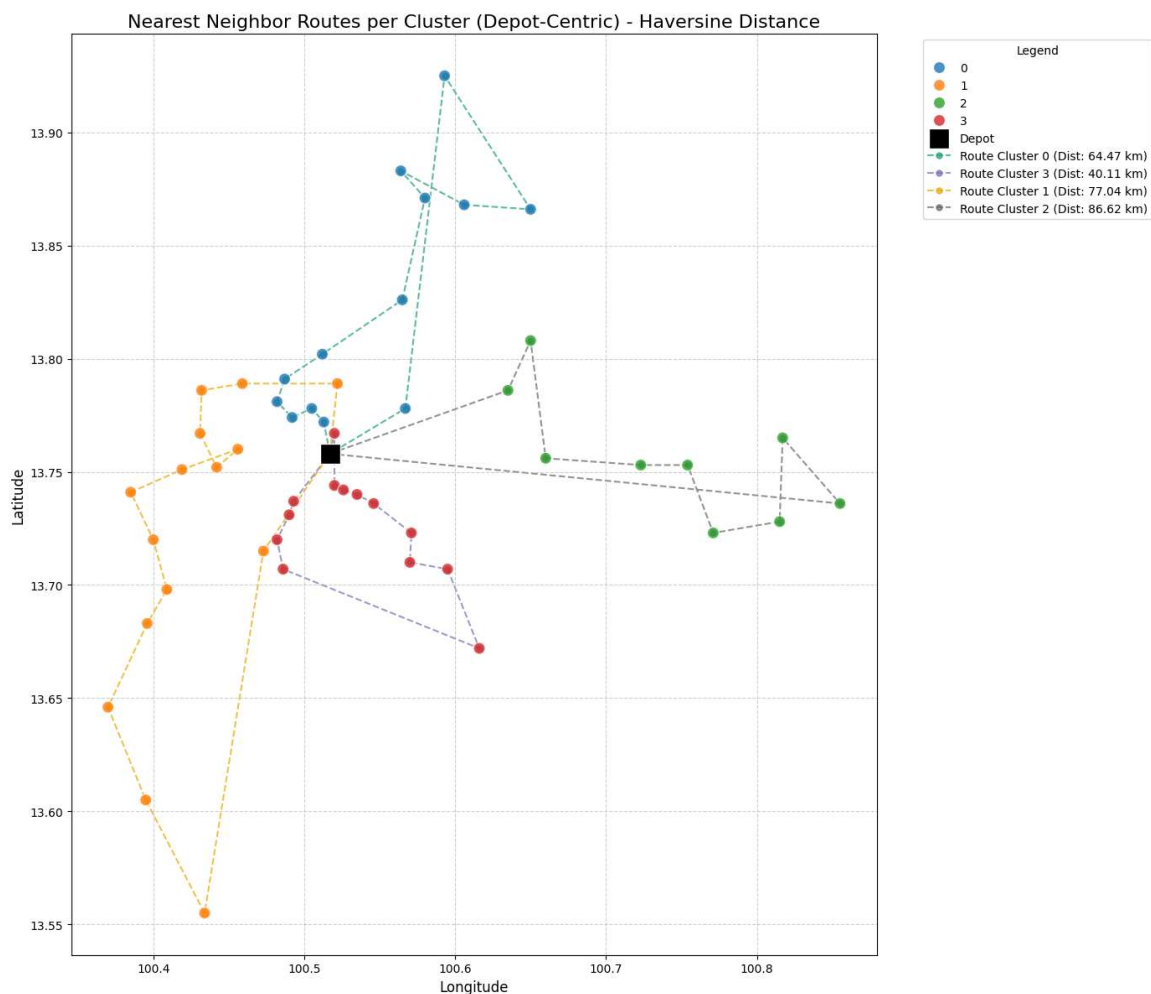
    # 3c.ii. Plot these coordinates as a line
    plt.plot(
        route_coords[:, 0],
        route_coords[:, 1],
        linestyle='--',
        marker='o',
        markersize=5,
        color=route_colors[i], # Use a distinct color for each route
        label=f'Route Cluster {cluster_id} (Dist: {total_route_distances[cluster_id]:.2f} km)',
        linewidth=1.5,
        alpha=0.7,
        zorder=1 # Ensure routes are behind points but visible
    )

# 3d. Add a title, 'Longitude' as the x-axis label, 'Latitude' as the y-axis label, and a clear
legend
plt.title('Nearest Neighbor Routes per Cluster (Depot-Centric) - Haversine Distance', fontsize=16)
plt.xlabel('Longitude', fontsize=12)
plt.ylabel('Latitude', fontsize=12)

# Create a combined legend
handles, labels = plt.gca().get_legend_handles_labels()
# Filter out duplicate legend entries if any (e.g., from seaborn's hue and explicit plot)
by_label = dict(zip(labels, handles))
plt.legend(by_label.values(), by_label.keys(), title='Legend', bbox_to_anchor=(1.05, 1),
loc='upper left')
```

บทที่ 11 การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อการจัดเส้นทางขนส่ง

```
plt.grid(True, linestyle='--', alpha=0.6)
plt.tight_layout()
plt.show()
```



ภาพที่ 11.3 การจัดเส้นทางขนส่งแบ่งเป็น 4 เส้นทาง ตามสีทั้ง 4 สี

ที่มา: ภาพโดยผู้แต่ง (ประมวลผลโดยโปรแกรม PYTHON)

บทที่ 11 การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อการจัดเส้นทางขนส่ง

เส้นทางของแต่ละ Cluster สามารถสรุปได้ดังต่อไปนี้

เส้นทางสำหรับ Cluster 0 (รวม 119 หน่วย)

จุดส่งของ	รหัส	เขต	ความต้องการ สินค้า
Point	Name	District	Demand
0	Depot	Dusit	0
1	DP13	Dusit	11
2	DP5	Dusit	7
3	DP15	Bang Phlat	10
4	DP14	Bang Phlat	8
5	DP17	Bang Phlat	13
6	DP16	Bang Phlat	4
7	DP45	Chatuchak	5
8	DP24	Lak Si	7
9	DP38	Lak Si	3
10	DP40	Bang Khen	15
11	DP41	Bang Khen	14
12	DP48	Don Mueang	13
13	DP1	Din Daeng	9
14	Depot	Dusit	0

บทที่ 11 การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อการจัดเส้นทางขนส่ง

เส้นทางสำหรับ Cluster 1 (รวม 95 หน่วย)

จุดส่งของ	รหัส	เขต	ความต้องการสินค้า
0	Depot	Dusit	0
1	DP23	Dusit	3
2	DP25	Taling Chan	2
3	DP27	Taling Chan	2
4	DP29	Taling Chan	5
5	DP22	Taling Chan	7
6	DP4	Taling Chan	9
7	DP20	Taling Chan	9
8	DP35	Bang Khae	4
9	DP32	Bang Khae	3
10	DP31	Bang Khae	13
11	DP28	Bang Khae	6
12	DP47	Bang Bon	7
13	DP49	Bang Khun Thian	13
14	DP50	Bang Khun Thian	10
15	DP9	Thon Buri	2
16	Depot	Dusit	0

บทที่ 11 การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อการจัดเส้นทางขนส่ง

เส้นทางสำหรับ Cluster 2 (รวม 67 หน่วย)

จุดส่งของ	รหัส	เขต	ความต้องการสินค้า
0	Depot	Dusit	0
1	DP33	Bang Kapi	7
2	DP42	Bung Kum	4
3	DP34	Bang Kapi	11
4	DP36	Lat Krabang	2
5	DP39	Lat Krabang	9
6	DP30	Lat Krabang	13
7	DP43	Lat Krabang	9
8	DP46	Lat Krabang	6
9	DP44	Lat Krabang	6
10	Depot	Dusit	0

บทที่ 11 การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อการจัดเส้นทางขนส่ง

เส้นทางสำหรับ Cluster 3 (รวม 106 หน่วย)

จุดส่งของ	รหัส	เขต	ความต้องการสินค้า
0	Depot	Dusit	0
1	DP10	Dusit	2
2	DP6	Pathum Wan	10
3	DP8	Pathum Wan	6
4	DP12	Pathum Wan	12
5	DP18	Pathum Wan	12
6	DP11	Khlong Toei	7
7	DP26	Khlong Toei	15
8	DP19	Khlong Toei	3
9	DP37	Bang Na	4
10	DP21	Thon Buri	8
11	DP7	Thon Buri	12
12	DP2	Thon Buri	9
13	DP3	Thon Buri	6
14	Depot	Dusit	0

11.5 สรุปผลลัพธ์จากการจัดเส้นทางขนส่ง

ผลการจัดเส้นทางขนส่งในกรณีศึกษาที่ใช้แนวทางแบบผสมระหว่าง SC และ NN ซึ่งสามารถแบ่งลูกค้าออกเป็น 4 กลุ่มตามความใกล้เคียงเชิงพื้นที่ และสร้างเส้นทางขนส่งในแต่ละกลุ่มได้อย่างเป็นระบบ โดยทุกเส้นทางเริ่มต้นและสิ้นสุดที่คลังสินค้า (depot) จากผลลัพธ์ของกรณีศึกษาพบว่าแต่ละ cluster มีปริมาณความต้องการสินค้าและเส้นทางที่แตกต่างกัน ดังนี้

1. Cluster 0: ความต้องการรวม 119 หน่วย ระยะทางประมาณ 64.41 km. ครอบคลุมพื้นที่ ดุสิต บางพลัด จตุจักร หลักสี่ บางเขน ดินแดง และดอนเมือง
2. Cluster 1: ความต้องการรวม 95 หน่วย ระยะทางประมาณ 40.11 km. ครอบคลุมพื้นที่ ดุสิต คลิ่งชัน บางแค บางบอน บางขุนเทียน และธนบุรี
3. Cluster 2: ความต้องการรวม 67 หน่วย ระยะทางประมาณ 86.62 km. ครอบคลุมพื้นที่ บางกะปิ ปิงกู๋ม และลาดกระบัง
4. Cluster 3: ความต้องการรวม 106 หน่วย ระยะทางประมาณ 77.04 km. ครอบคลุมพื้นที่ ดุสิต ปทุมวัน คลองเตย บางนา และธนบุรี

เมื่อเปรียบเทียบกับข้อจำกัดของระบบที่มีรถขนส่ง 4 คัน และความจุคันละ 120 หน่วย ระยะทางรวม 268 km. พบว่าทุก Cluster มีความต้องการไม่เกินความจุของรถในแต่ละคัน ซึ่งหมายความว่าสามารถให้บริการได้จริงโดยใช้รถ 1 คันต่อ 1 Cluster

นอกจากนี้ จากแผนภาพเส้นทาง แสดงให้เห็นว่าเส้นทางในแต่ละ Cluster มีลักษณะเป็นกลุ่มก้อน (geographical grouping) และลดการวิ่งข้ามพื้นที่ที่ไม่จำเป็น ส่งผลให้ระยะทางรวมของแต่ละเส้นทางมีประสิทธิภาพ และช่วยลดต้นทุนการขนส่งได้อย่างมีนัยสำคัญ โดยสรุป วิธีการผสมระหว่าง SC และ NN สามารถจัดเส้นทางขนส่งได้อย่างมีประสิทธิภาพ ทั้งในด้านการแบ่งพื้นที่ให้บริการ การควบคุมข้อจำกัดด้านความจุ และการลดระยะทางรวมของระบบ ซึ่งเหมาะสมสำหรับการประยุกต์ใช้ในงานโลจิสติกส์จริง และสามารถต่อยอดไปสู่การใช้เทคนิคขั้นสูงเพื่อเพิ่มประสิทธิภาพได้ในอนาคต

เอกสารอ้างอิง

- Areerakulkan, N., Lhaochot, L., & Kijwattanaphokin, S. (2026). Vehicle Routing Planning Using a Hybrid Approach of Spectral Clustering and Nearest Neighbour Algorithms. *Journal of the CUFST*, 1(1).
- Kool, W., van Hoof, H., & Welling, M. (2019). *Attention, learn to solve routing problems!* arXiv. <https://doi.org/10.48550/arXiv.1803.08475>.
- Laporte, G. (2009). Fifty years of vehicle routing. *Transportation Science*, 43(4), 408-416. <https://doi.org/10.1287/trsc.1090.0301>.
- Nazari, M., Oroojlooy, A., Snyder, L. V., & Takáč, M. (2018). *Reinforcement learning for solving the vehicle routing problem.* arXiv. <https://doi.org/10.48550/arXiv.1802.04240>.
- Shi, J., & Malik, J. (2000). Normalized cuts and image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22(8), 888-905. <https://doi.org/10.1109/34.868688>.
- Toth, P., & Vigo, D. (2014). *Vehicle Routing: Problems, Methods, and Applications.* Society for Industrial and Applied Mathematics. <https://doi.org/10.1137/1.9781611973594>.
- Von Luxburg, U. (2007). A tutorial on spectral clustering. *Statistics and Computing*, 17(4), 395-416. <https://doi.org/10.1007/s11222-007-9033-z>.
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77-84. <https://doi.org/10.1111/jbl.12010>.

บทที่ 12

งานวิจัยที่น่าสนใจเกี่ยวกับการคัดเลือกซัพพลายเออร์อย่างยั่งยืนด้วย กระบวนการวิเคราะห์เชิงลำดับชั้น

Interesting Research on Sustainable Supplier Selection Using the Analytic Hierarchy Process

หัวข้อ

- 12.1 วัตถุประสงค์ของงานวิจัย
- 12.2 บริบทธุรกิจ
- 12.3 ชุดข้อมูลและตัวแปรที่ใช้
- 12.4 ขั้นตอนการวิเคราะห์ด้วยวิธี AHP
- 12.5 สรุปผลงานวิจัย

ในปัจจุบัน ภาคอุตสาหกรรมและธุรกิจมีบทบาทสำคัญต่อการขับเคลื่อนเศรษฐกิจของประเทศ แต่การเติบโตดังกล่าวก่อให้เกิดผลกระทบด้านสิ่งแวดล้อมและสังคมอย่างหลีกเลี่ยงไม่ได้ ส่งผลให้องค์กรต่าง ๆ ให้ความสำคัญกับแนวคิด การพัฒนาอย่างยั่งยืน (sustainable development) ซึ่งครอบคลุม 3 มิติหลัก ได้แก่ ด้านเศรษฐกิจ สังคม และสิ่งแวดล้อม โดยแนวคิดดังกล่าวได้รับการพัฒนาและยอมรับในระดับสากลผ่านรายงาน Brundtland (Holden et al., 2014) และได้รับการอธิบายเพิ่มเติมในเชิงโครงสร้างของความยั่งยืนในงานของ Purvis et al. (2019) ในบริบทของการจัดการซัพพลายเชน กระบวนการจัดซื้อและการคัดเลือกผู้ขาย (supplier selection) ถือเป็นกิจกรรมที่มีความสำคัญอย่างยิ่ง เนื่องจากเป็นจุดเริ่มต้นของการไหลของสินค้าและวัตถุดิบเข้าสู่ระบบการผลิตและการกระจายสินค้า โดยในอดีต องค์กรมักพิจารณาผู้ขายจากปัจจัยด้านราคาและการส่งมอบเป็นหลัก แต่ในปัจจุบัน แนวคิด การคัดเลือกผู้ขายอย่างยั่งยืน (sustainable supplier selection, SSS) ได้เข้ามามีบทบาทมากขึ้น โดยต้องพิจารณาปัจจัยด้านสิ่งแวดล้อมและสังคมร่วมด้วย (Chaitongrat & Areerakulkan, 2021) การคัดเลือกผู้ขายอย่างยั่งยืนเป็นปัญหาที่มีความ

ซับซ้อน เนื่องจากต้องพิจารณาหลายเกณฑ์ที่มีความสัมพันธ์กันทั้งเชิงปริมาณและเชิงคุณภาพ จึงจัดอยู่ในกลุ่มปัญหาการตัดสินใจหลายเกณฑ์ (multi-criteria decision making, MCDM) ซึ่งจำเป็นต้องใช้เครื่องมือวิเคราะห์ที่มีประสิทธิภาพ หนึ่งในวิธีที่ได้รับความนิยมอย่างแพร่หลายคือ กระบวนการวิเคราะห์เชิงลำดับชั้น (analytic hierarchy process, AHP) ซึ่งถูกพัฒนาโดย Saaty (1980) และสามารถใช้ในการจัดลำดับความสำคัญของเกณฑ์และทางเลือกได้อย่างเป็นระบบ โดยงานวิจัยของ Chaitongrat และ Areerakulkan (2021) ได้แสดงให้เห็นว่า AHP สามารถนำมาใช้ประเมินและคัดเลือกผู้ขายภายใต้บริบทของความยั่งยืนได้อย่างมีประสิทธิภาพ

จากการทบทวนวรรณกรรม พบว่างานวิจัยด้านการคัดเลือกผู้ขายอย่างยั่งยืนส่วนใหญ่ใช้เกณฑ์การประเมินที่สอดคล้องกับแนวคิด Triple Bottom Line (TPL) ได้แก่ มิติด้านเศรษฐกิจ สิ่งแวดล้อม และสังคม (Memari et al., 2019; Pishchulov et al., 2019) ซึ่งสอดคล้องกับผลการศึกษาของ Chaitongrat และ Areerakulkan (2021) ที่พบว่าเกณฑ์ด้านเศรษฐกิจยังคงมีความสำคัญสูงสุด รองลงมาคือด้านสิ่งแวดล้อมและสังคม นอกจากนี้ ยังมีการพัฒนาเทคนิคการตัดสินใจขั้นสูง เช่น Fuzzy AHP TOPSIS และ DEA เพื่อเพิ่มความแม่นยำในการประเมินผู้ขายในสภาพแวดล้อมที่มีความไม่แน่นอน (Amindoust et al., 2012)

ดังนั้น ในบทนี้จึงมุ่งเน้นอธิบายถึงวิธีการคัดเลือกผู้ขายอย่างยั่งยืนด้วยวิธี AHP โดยใช้งานวิจัยกรณีศึกษาธุรกิจค้าปลีก เพื่อพัฒนาแนวทางการตัดสินใจที่เป็นระบบ สามารถสะท้อนความสำคัญของเกณฑ์ต่าง ๆ ได้อย่างชัดเจน และสนับสนุนการบริหารจัดการซัพพลายเชนให้มีประสิทธิภาพและยั่งยืนในระยะยาว

12.1 วัตถุประสงค์ของงานวิจัย

งานวิจัยนี้อ้างอิงจากการศึกษาของ Areerakulkan et al. (2021) มีวัตถุประสงค์เพื่อพัฒนาแนวทางการคัดเลือกผู้ขายอย่างยั่งยืนในบริบทของธุรกิจค้าปลีก โดยมุ่งเน้นการประยุกต์ใช้กระบวนการวิเคราะห์เชิงลำดับชั้น (Analytic Hierarchy Process, AHP) ซึ่งเป็นเครื่องมือที่มีประสิทธิภาพในการตัดสินใจพหุเกณฑ์ที่มีความซับซ้อน ทั้งนี้ การศึกษามุ่งเน้นการกำหนดเกณฑ์การประเมินผู้ขายที่ครอบคลุมมิติด้านเศรษฐกิจ สิ่งแวดล้อม และสังคม ตามแนวคิดการพัฒนาอย่างยั่งยืน เพื่อให้การตัดสินใจมีความรอบด้านและสอดคล้องกับแนวโน้มของการบริหารจัดการซัพพลายเชนในปัจจุบัน นอกจากนี้ งานวิจัยยังมีวัตถุประสงค์เพื่อพัฒนาแบบจำลองการตัดสินใจที่สามารถกำหนดค่า

น้ำหนักความสำคัญของเกณฑ์หลักและเกณฑ์ย่อยได้อย่างเป็นระบบ โดยอาศัยการเปรียบเทียบเชิงคู่ (pairwise comparison) เพื่อสะท้อนความสำคัญของปัจจัยต่าง ๆ ที่มีผลต่อการคัดเลือกผู้ขาย อีกทั้งยังมุ่งประเมินและจัดอันดับผู้ขายที่เหมาะสมจากกลุ่มผู้ขายที่มีอยู่ในกรณีศึกษา โดยใช้ข้อมูลเชิงประจักษ์จากธุรกิจค้าปลีกที่มีผู้ขายหลายราย นอกจากนี้การศึกษานี้มีเป้าหมายเพื่อสนับสนุนการตัดสินใจเชิงกลยุทธ์ด้านการจัดซื้อและการบริหารซัพพลายเชน โดยเน้นการเพิ่มประสิทธิภาพในการดำเนินงาน ควบคู่กับการส่งเสริมความยั่งยืนขององค์กรในระยะยาว ซึ่งผลลัพธ์ที่ได้สามารถนำไปใช้เป็นแนวทางหรือเครื่องมือในการประเมินและคัดเลือกผู้ขายในสถานการณ์จริงได้อย่างมีประสิทธิภาพและเป็นระบบ

12.2 บริบทธุรกิจ

งานวิจัยนี้อยู่ในบริบทของธุรกิจค้าปลีกที่ดำเนินการจัดจำหน่ายสินค้าอุปโภคบริโภค โดยเฉพาะสินค้าประเภทอาหารสำเร็จรูปและเครื่องดื่ม ซึ่งมีความต้องการของตลาดที่ผันผวนและต้องการความรวดเร็วในการตอบสนองต่อผู้บริโภค ธุรกิจลักษณะนี้มีการดำเนินงานผ่านเครือข่ายสาขาที่กระจายอยู่ในหลายพื้นที่ ทำให้ต้องพึ่งพาระบบซัพพลายเชนที่มีประสิทธิภาพเพื่อให้สามารถจัดส่งสินค้าได้อย่างต่อเนื่องและเพียงพอต่อความต้องการของลูกค้า ในกระบวนการจัดซื้อ องค์กรจำเป็นต้องพึ่งพาผู้ขายหลักหลายรายที่จัดจำหน่ายสินค้าประเภทเดียวกัน โดยในกรณีศึกษานี้มีผู้ขายหลักจำนวน 4 ราย ซึ่งแต่ละรายมีศักยภาพและจุดเด่นที่แตกต่างกัน ทั้งในด้านต้นทุน คุณภาพสินค้า ความสามารถในการส่งมอบ และการให้บริการ อย่างไรก็ตาม ในอดีตองค์กรยังคงใช้เกณฑ์การคัดเลือกผู้ขายที่เน้นเพียงด้านราคาและการส่งมอบเป็นหลัก ซึ่งอาจไม่เพียงพอในการรองรับแนวโน้มการแข่งขันและความต้องการด้านความยั่งยืนในปัจจุบัน

ในขณะเดียวกัน องค์กรมีวิสัยทัศน์ในการพัฒนาไปสู่การเป็นองค์กรที่มีความยั่งยืนและมีความรับผิดชอบต่อสังคม แต่ยังคงขาดกรอบแนวทางและเครื่องมือที่ชัดเจนในการคัดเลือกและประเมินผู้ขายให้สอดคล้องกับแนวคิดดังกล่าว ส่งผลให้เกิดความจำเป็นในการพัฒนาแนวทาง การตัดสินใจที่สามารถพิจารณาปัจจัยได้อย่างรอบด้าน ครอบคลุมทั้งมิติด้านเศรษฐกิจ สิ่งแวดล้อม และสังคม ดังนั้นบริบทของธุรกิจในงานวิจัยนี้จึงสะท้อนให้เห็นถึงความท้าทายขององค์กรค้าปลีกในยุคปัจจุบัน ที่ต้องเผชิญกับการแข่งขันสูง ความไม่แน่นอนของอุปสงค์ และแรงกดดันด้านความยั่งยืน ซึ่งล้วนเป็นปัจจัยสำคัญที่ทำให้การคัดเลือกผู้ขายไม่สามารถพิจารณาเพียงมิติเดียวได้อีกต่อไป แต่จำเป็นต้องใช้แนวทาง การตัดสินใจเชิงวิเคราะห์ที่เป็นระบบ เพื่อเพิ่มประสิทธิภาพของซัพพลายเชนและสร้างความได้เปรียบทางการแข่งขันในระยะยาว

12.3 ชุดข้อมูลและตัวแปรที่ใช้

ในการศึกษานี้ ชุดข้อมูลที่ใช้เป็นข้อมูลเชิงประเมิน (evaluation data) ที่ได้จากการรวบรวมความคิดเห็นของผู้เชี่ยวชาญภายในองค์กร ซึ่งมีบทบาทโดยตรงในกระบวนการจัดซื้อและการบริหารคุณภาพ โดยประกอบด้วยผู้เชี่ยวชาญจำนวน 4 คน ได้แก่ ผู้บริหารระดับสูง 2 คน ผู้จัดการฝ่ายจัดซื้อ 1 คน และผู้จัดการฝ่ายประกันคุณภาพ 1 คน ซึ่งล้วนเป็นผู้ที่มีประสบการณ์และความรู้เกี่ยวกับการคัดเลือกผู้ขายและการตัดสินใจเชิงกลยุทธ์ขององค์กร ผู้เชี่ยวชาญเหล่านี้มีหน้าที่ในการให้คะแนนความสำคัญของเกณฑ์และเกณฑ์ย่อย รวมถึงการเปรียบเทียบเชิงคู่ (pairwise comparison) ตามกระบวนการวิเคราะห์เชิงลำดับชั้น (AHP) เพื่อสะท้อนมุมมองเชิงประสบการณ์และความสำคัญของปัจจัยต่าง ๆ ในบริบทจริงขององค์กร นอกจากนี้ ยังมีการใช้วิธีการระดมความคิดเห็น (brainstorming) เพื่อคัดเลือกตัวแปรที่เหมาะสม โดยกำหนดให้เกณฑ์ที่ได้รับการยอมรับจากผู้เชี่ยวชาญตั้งแต่ 3 คนขึ้นไปถูกนำมาใช้ในการวิเคราะห์ เพื่อเพิ่มความน่าเชื่อถือของข้อมูล

ชุดข้อมูลประกอบด้วยผู้ขายหลักจำนวน 4 ราย ซึ่งเป็นผู้จัดจำหน่ายสินค้าประเภทเดียวกันในธุรกิจค้าปลีก โดยผู้ขายแต่ละรายจะถูกประเมินภายใต้เกณฑ์การคัดเลือกผู้ขายอย่างยั่งยืนที่พัฒนาขึ้นจากการทบทวนวรรณกรรมและการระดมความคิดเห็นของผู้เชี่ยวชาญ

ตัวแปรที่ใช้ในการวิเคราะห์สามารถแบ่งออกเป็น 3 ระดับ ได้แก่

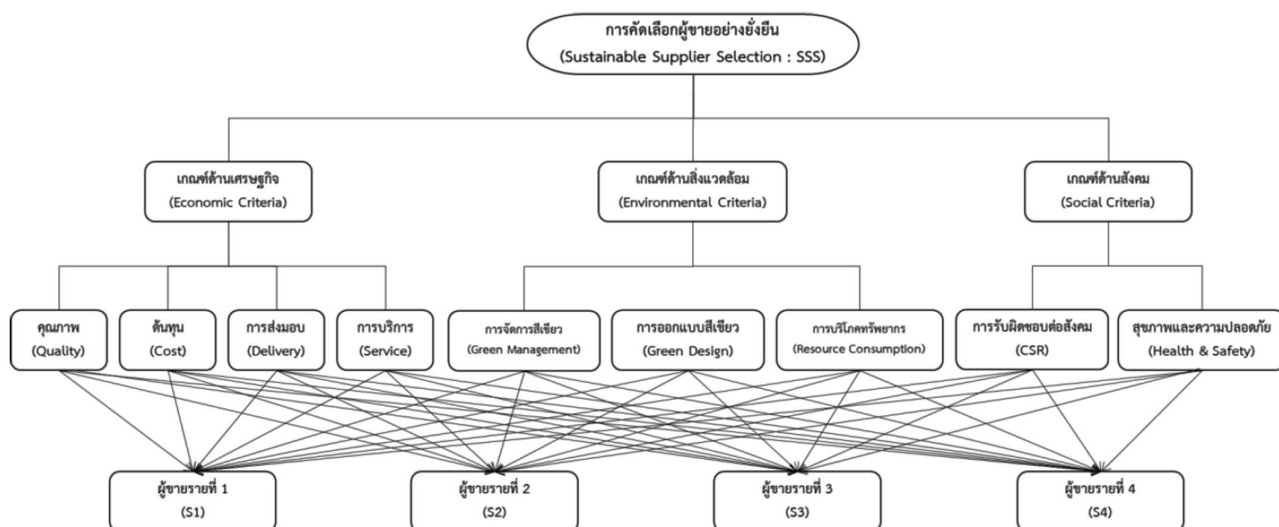
1. ตัวแปรระดับเกณฑ์หลัก ประกอบด้วย 3 มิติหลักตามแนวคิดความยั่งยืน ได้แก่ มิติด้านเศรษฐกิจ (economic) มิติด้านสิ่งแวดล้อม (environmental) และมิติด้านสังคม (social)
2. ตัวแปรระดับเกณฑ์ย่อยในแต่ละเกณฑ์หลักประกอบด้วยตัวแปรย่อย ดังนี้ ด้านเศรษฐกิจ ได้แก่ คุณภาพสินค้า/บริการ (quality) ต้นทุน (cost) การส่งมอบ (delivery) และการบริการ (service) ด้านสิ่งแวดล้อม ได้แก่ การจัดการสิ่งแวดล้อม (green management) การออกแบบเพื่อสิ่งแวดล้อม (eco-design) และการใช้ทรัพยากรและพลังงาน (resource & energy consumption) ด้านสังคม ได้แก่ ความรับผิดชอบต่อสังคม (CSR) และ สุขภาพและความปลอดภัย (health & safety) ตัวแปรย่อยทั้งหมดนี้ได้มาจากการคัดเลือกผ่านกระบวนการระดมความคิดเห็น (brainstorming) และการโหวตจากผู้เชี่ยวชาญ เพื่อให้ได้ตัวแปรที่เหมาะสมกับบริบทขององค์กร
3. ตัวแปรทางเลือก (alternatives) ประกอบด้วยผู้ขายจำนวน 4 ราย ได้แก่ ผู้ขาย A B C และ D ซึ่งเป็นตัวเลือกที่ต้องนำมาประเมินและจัดอันดับ โดยข้อมูลทั้งหมดจะถูกนำเข้าสู่วิธีการ AHP เพื่อคำนวณค่าน้ำหนักความสำคัญของแต่ละเกณฑ์ และใช้ในการจัดอันดับผู้ขายที่เหมาะสมที่สุดภายใต้กรอบแนวคิดความยั่งยืน

12.4 ขั้นตอนการวิเคราะห์ด้วยวิธี AHP

การวิเคราะห์ข้อมูลในงานวิจัยนี้ใช้กระบวนการวิเคราะห์เชิงลำดับชั้น (Analytic Hierarchy Process, AHP) ซึ่งเป็นวิธีการตัดสินใจแบบพหุเกณฑ์ที่ช่วยจัดลำดับความสำคัญของปัจจัยต่าง ๆ อย่างเป็นระบบ โดยสามารถอธิบายขั้นตอนการดำเนินการได้ดังนี้

ขั้นตอนที่ 1 การกำหนดปัญหาและวัตถุประสงค์ เริ่มต้นจากการกำหนดวัตถุประสงค์ของการศึกษา คือ การคัดเลือกผู้ขายที่เหมาะสมภายใต้แนวคิดความยั่งยืน โดยพิจารณาผู้ขายจำนวน 4 ราย และกำหนดกรอบการตัดสินใจให้ชัดเจน

ขั้นตอนที่ 2 การกำหนดเกณฑ์และโครงสร้างลำดับชั้น กำหนดเกณฑ์การตัดสินใจออกเป็นลำดับชั้น (hierarchy structure) ประกอบด้วย 3 ระดับ ได้แก่ ระดับเป้าหมาย (goal) ระดับเกณฑ์หลัก (criteria) และระดับเกณฑ์ย่อย (sub-criteria) รวมถึงระดับทางเลือก (alternatives) ซึ่งในงานวิจัยนี้ใช้เกณฑ์หลัก 3 ด้าน คือ เศรษฐกิจ สิ่งแวดล้อม และสังคม โดยแผนภูมิลำดับชั้นเชิงวิเคราะห์แสดงดังภาพที่ 12.1



ภาพที่ 12.1 แผนภูมิลำดับชั้นเชิงวิเคราะห์ในการคัดเลือกผู้ขายอย่างยั่งยืน
ที่มา: (Chaitongrat & Areerakulkan, 2021)

ขั้นตอนที่ 3 การสร้างเมทริกซ์เปรียบเทียบเชิงคู่ (Pairwise Comparison Matrix) ทำการเปรียบเทียบความสำคัญของเกณฑ์และเกณฑ์ย่อยแบบเป็นคู่ ๆ โดยใช้มาตราส่วน 1-9 ของ Saaty เพื่อสะท้อนระดับความสำคัญของแต่ละปัจจัย จากความคิดเห็นของผู้เชี่ยวชาญ

ให้ A เป็นเมทริกซ์เปรียบเทียบเชิงคู่ของเกณฑ์จำนวน n ตัว จะได้ว่า

$$A = [a_{ij}]_{n \times n} \quad 12.1$$

โดยที่

$$A = \begin{bmatrix} 1 & a_{12} & a_{13} & \cdots & a_{1n} \\ \frac{1}{a_{12}} & 1 & a_{23} & \cdots & a_{2n} \\ \frac{1}{a_{13}} & \frac{1}{a_{23}} & 1 & \cdots & a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{1}{a_{1n}} & \frac{1}{a_{2n}} & \frac{1}{a_{3n}} & \cdots & 1 \end{bmatrix} \quad 12.2$$

เมื่อ a_{ij} คือค่าความสำคัญของเกณฑ์ i เมื่อเทียบกับเกณฑ์ j ตามมาตราส่วน 1-9 ของ Saaty และมีคุณสมบัติว่า

$$a_{ji} = \frac{1}{a_{ij}}, a_{ii} = 1$$

ตัวอย่าง 12.1 การสร้างเมทริกซ์เปรียบเทียบเชิงคู่

กรณีศึกษาในงานวิจัยของ (Chaitongrat & Areerakulkan, 2021) กำหนดเกณฑ์หลักและเกณฑ์รองแสดงดังภาพที่ 12.1 เพื่อให้เข้าใจง่าย จะยกตัวอย่างการหาน้ำหนักของ เกณฑ์หลัก 3 ด้าน ได้แก่

- C_1 = ด้านเศรษฐกิจ
- C_2 = ด้านสิ่งแวดล้อม
- C_3 = ด้านสังคม

สมมติว่าผู้เชี่ยวชาญให้ความเห็นดังนี้

- ด้านเศรษฐกิจสำคัญกว่าด้านสิ่งแวดล้อม 5 เท่า
- ด้านเศรษฐกิจสำคัญกว่าด้านสังคม 7 เท่า
- ด้านสิ่งแวดล้อมสำคัญกว่าด้านสังคม 3 เท่า

ดังนั้น เมทริกซ์เปรียบเทียบเชิงคู่คือ

บทที่ 12 งานวิจัยที่น่าสนใจเกี่ยวกับการคัดเลือกซัพพลายเออร์อย่างยั่งยืน
ด้วยกระบวนการวิเคราะห์เชิงลำดับชั้น

$$A = \begin{bmatrix} 1 & 5 & 7 \\ \frac{1}{5} & 1 & 3 \\ \frac{1}{7} & \frac{1}{3} & 1 \end{bmatrix}$$

หรือในรูปทศนิยม

$$A \approx \begin{bmatrix} 1 & 5 & 7 \\ 0.2 & 1 & 3 \\ 0.1429 & 0.3333 & 1 \end{bmatrix}$$

ขั้นตอนที่ 4 การคำนวณค่าน้ำหนัก (Weight Calculation) นำค่าที่ได้จากเมทริกซ์มาเข้าสู่กระบวนการคำนวณหาค่าน้ำหนักของแต่ละเกณฑ์ โดยใช้วิธีค่าเฉลี่ยเรขาคณิต (Geometric Mean) และทำการปรับให้อยู่ในรูป Normalized Weight เพื่อให้ผลรวมของน้ำหนักเท่ากับ 1 สำหรับแถวที่ i คำนวณค่าเฉลี่ยเรขาคณิตได้เป็น

$$V_i = \left(\prod_{j=1}^n a_{ij} \right)^{1/n} \quad 12.3$$

เมื่อได้ V_i ของทุกแถวแล้ว จึงนอร์มัลไลซ์เพื่อหาน้ำหนัก

$$w_i = \frac{V_i}{\sum_{i=1}^n V_i} \quad 12.4$$

หรือเขียนในรูปเวกเตอร์ได้เป็น

$$\mathbf{w} = \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_n \end{bmatrix} \quad 12.5$$

โดยที่

$$\sum_{i=1}^n w_i = 1$$

ตัวอย่าง 12.2 การสร้างเมตริกซ์เปรียบเทียบเชิงคู่

ขั้นที่ 1 หาค่าเฉลี่ยเรขาคณิตของแต่ละแถว ใช้ข้อมูลกรณีศึกษาในงานวิจัยของ (Chaitongrat & Areerakulkan, 2021)

แถวที่ 1: ด้านเศรษฐกิจ

$$V_1 = (1 \times 5 \times 7)^{1/3} = 35^{1/3} \approx 3.2711$$

แถวที่ 2: ด้านสิ่งแวดล้อม

$$V_2 = \left(\frac{1}{5} \times 1 \times 3\right)^{1/3} = 0.6^{1/3} \approx 0.8434$$

แถวที่ 3: ด้านสังคม

$$V_3 = \left(\frac{1}{7} \times \frac{1}{3} \times 1\right)^{1/3} = \left(\frac{1}{21}\right)^{1/3} \approx 0.3620$$

ดังนั้น เวกเตอร์ค่าเฉลี่ยเรขาคณิตคือ

$$\mathbf{V} = \begin{bmatrix} 3.2711 \\ 0.8434 \\ 0.3620 \end{bmatrix}$$

ขั้นที่ 2 นอร์มัลไลซ์เพื่อหาน้ำหนัก

ผลรวมของ V_i เท่ากับ

$$\sum V_i = 3.2711 + 0.8434 + 0.3620 = 4.4765$$

ดังนั้น

$$w_1 = \frac{3.2711}{4.4765} \approx 0.7307$$
$$w_2 = \frac{0.8434}{4.4765} \approx 0.1884$$

บทที่ 12 งานวิจัยที่น่าสนใจเกี่ยวกับการคัดเลือกซ์พปลายเออร์อย่างยั่งยืน
ด้วยกระบวนการวิเคราะห์เชิงลำดับชั้น

$$w_3 = \frac{0.3620}{4.4765} \approx 0.0809$$

ดังนั้น เวกเตอร์น้ำหนักคือ

$$\mathbf{w} = \begin{bmatrix} 0.7307 \\ 0.1884 \\ 0.0809 \end{bmatrix}$$

หรือประมาณได้เป็น

ด้านเศรษฐกิจ = 0.731, ด้านสิ่งแวดล้อม = 0.188, ด้านสังคม = 0.081

ขั้นตอนที่ 5 การตรวจสอบความสอดคล้อง (Consistency Check) ตรวจสอบความสอดคล้องของการเปรียบเทียบด้วยค่า Consistency Ratio (C.R.) โดยค่าที่เหมาะสมต้องน้อยกว่า 0.10 เพื่อให้มั่นใจว่าการให้คะแนนของผู้เชี่ยวชาญมีความสมเหตุสมผล

สำหรับขั้นตอนการตรวจสอบความสอดคล้องจะเริ่มจากการหาค่า $\lambda_{\max}$ โดยเมื่อได้เวกเตอร์น้ำหนัก $\mathbf{w}$ แล้ว คำนวณ $A\mathbf{w}$ จากนั้นหาค่า

$$\lambda_i = \frac{(A\mathbf{w})_i}{w_i} \quad 12.6$$

แล้วเฉลี่ยเป็น

$$\lambda_{\max} = \frac{1}{n} \sum_{i=1}^n \lambda_i \quad 12.7$$

5) การตรวจสอบความสอดคล้อง

คำนวณดัชนีความสอดคล้อง

$$CI = \frac{\lambda_{\max} - n}{n - 1} \quad 12.8$$

และอัตราส่วนความสอดคล้อง

$$CR = \frac{CI}{RI} \quad 12.9$$

โดย RI คือ Random Index ซึ่งขึ้นกับจำนวนเกณฑ์ n หาก $CR < 0.10$ ถือว่าการเปรียบเทียบมีความสอดคล้องเพียงพอ ซึ่งไฟล์กรณีศึกษานี้รายงานค่า overall consistency ratio เท่ากับ 0.07 จึงถือว่ายอมรับได้

ตัวอย่าง 12.3 การตรวจสอบคะแนนความสอดคล้องของการประเมิน

ใช้ข้อมูลกรณีศึกษาในงานวิจัยของ (Chaitongrat & Areerakulkan, 2021) จากตัวอย่าง 12.1-12.3 ขั้นตอนการคำนวณค่าคะแนนความสอดคล้องใช้สมการที่ 12.6-12.9 ในขั้นตอนที่ 5 แสดงดังต่อไปนี้

1. คำนวณ A_w

$$A_w = \begin{bmatrix} 1 & 5 & 7 \\ 0.2 & 1 & 3 \\ 0.1429 & 0.3333 & 1 \end{bmatrix} \begin{bmatrix} 0.7307 \\ 0.1884 \\ 0.0809 \end{bmatrix}$$

คำนวณทีละแถวได้ดังนี้

แถวที่ 1

$$(1)(0.7307) + (5)(0.1884) + (7)(0.0809) = 0.7307 + 0.9420 + 0.5663 = 2.2390$$

แถวที่ 2

$$(0.2)(0.7307) + (1)(0.1884) + (3)(0.0809) = 0.1461 + 0.1884 + 0.2427 = 0.5772$$

แถวที่ 3

$$(0.1429)(0.7307) + (0.3333)(0.1884) + (1)(0.0809) \approx 0.1044 + 0.0628 + 0.0809 = 0.2481$$

ดังนั้น

$$A_w = \begin{bmatrix} 2.2390 \\ 0.5772 \\ 0.2481 \end{bmatrix}$$

2. หา λ_i และ $\lambda_{\max}$

$$\lambda_1 = \frac{2.2390}{0.7307} \approx 3.0638$$

บทที่ 12 งานวิจัยที่น่าสนใจเกี่ยวกับการคัดเลือกซัพพลายเออร์อย่างยั่งยืน
ด้วยกระบวนการวิเคราะห์เชิงลำดับชั้น

$$\lambda_2 = \frac{0.5772}{0.1884} \approx 3.0637$$

$$\lambda_3 = \frac{0.2481}{0.0809} \approx 3.0674$$

ดังนั้น

$$\lambda_{\max} = \frac{3.0638 + 3.0637 + 3.0674}{3} \approx 3.0649$$

3. หา CI และ CR

เมื่อ $n = 3$

$$CI = \frac{\lambda_{\max} - n}{n - 1} = \frac{3.0649 - 3}{2} = \frac{0.0649}{2} = 0.03245$$

สำหรับ $n = 3$ ค่า $RI = 0.58$

ดังนั้น

$$CR = \frac{CI}{RI} = \frac{0.03245}{0.58} \approx 0.0560$$

เนื่องจาก

$$CR = 0.056 < 0.10$$

จึงสรุปได้ว่าเมทริกซ์นี้มีความสอดคล้องและสามารถใช้ค่าน้ำหนักที่คำนวณได้ในการตัดสินใจ

ขั้นตอนที่ 6 การคำนวณคะแนนรวมและจัดอันดับทางเลือก นำค่าน้ำหนักของเกณฑ์และคะแนนของผู้ขายแต่ละรายมาคำนวณเพื่อหาคะแนนรวม (overall score) และจัดอันดับผู้ขายที่เหมาะสมที่สุด โดยผู้ขายที่มีค่าน้ำหนักรวมสูงสุดจะถูกเลือกเป็นทางเลือกที่ดีที่สุด

เมื่อได้ค่าน้ำหนักของเกณฑ์แล้ว ขั้นตอนต่อไปคือการคำนวณคะแนนรวมของผู้ขายแต่ละราย โดยใช้หลักการถ่วงน้ำหนัก ใช้ข้อมูลจากตัวอย่างที่ผ่านมาจะได้ว่า

$$\mathbf{w} \approx \begin{bmatrix} 0.731 \\ 0.188 \\ 0.081 \end{bmatrix}$$

บทที่ 12 งานวิจัยที่น่าสนใจเกี่ยวกับการคัดเลือกซ์พพลายเออร์อย่างยั่งยืน
ด้วยกระบวนการวิเคราะห์เชิงลำดับชั้น

และสมมติว่าผู้ชาย B มีคะแนนภายใต้ 3 เกณฑ์หลักดังนี้

- ด้านเศรษฐกิจ = 0.40
- ด้านสิ่งแวดล้อม = 0.25
- ด้านสังคม = 0.20

คะแนนรวมของผู้ชาย B คือ

$$Score_B = (0.731)(0.40) + (0.188)(0.25) + (0.081)(0.20)$$

$$Score_B = 0.2924 + 0.0470 + 0.0162 = 0.3556$$

ดังนั้นคะแนนรวมประมาณเท่ากับ

$$Score_B \approx 0.356$$

ซึ่งมีค่าใกล้เคียงกับผลในกรณีศึกษาที่ผู้ชาย B มีค่าน้ำหนักรวมประมาณ 0.357 และได้รับการจัดอันดับเป็นอันดับ 1

12.5 สรุปผลงานวิจัย

กรณีศึกษาของงานวิจัยที่กำลังนำเสนอ ใช้การพิจารณาองค์ประกอบในการตัดสินใจ ได้มาจากการทบทวนทวนวรรณกรรมและการระดมสมองของผู้เชี่ยวชาญ เพื่อนำมาสร้างเป็นแผนภูมิลำดับชั้นดังที่แสดงในภาพที่ 12.1 การวินิจฉัยเปรียบเทียบความสำคัญของเกณฑ์ในการตัดสินใจ โดยผู้เชี่ยวชาญได้ผลการวิเคราะห์น้ำหนักเกณฑ์หลักและเกณฑ์รองเพื่อทำการประเมินและคัดเลือกผู้ขาย ด้วยเกณฑ์ความยั่งยืน ผลของการวิเคราะห์น้ำหนักแสดงดังตารางที่ 12.1-12.2

ตารางที่ 12.1 ผลการวิเคราะห์น้ำหนักเกณฑ์หลักและเกณฑ์รอง

เกณฑ์	ค่าเฉลี่ย เกณฑ์หลัก	ค่าเฉลี่ย เกณฑ์รอง
ด้านเศรษฐกิจ (Economic Criteria)	0.731	
- ต้นทุน (Cost)		0.574
- คุณภาพสินค้าหรือบริการ (Quality)		0.282
- การบริการ (Service)		0.092
- การส่งมอบ (Delivery)		0.052
ด้านสิ่งแวดล้อม (Environmental Criteria)	0.188	
- การบริโภคพลังงานและทรัพยากร (Resource and Energy consumption)		0.649
- การจัดการสีเขียว (Green/Environmental Management)		0.279
- การออกแบบเพื่อสิ่งแวดล้อม (Green Design/Eco Design)		0.072
ด้านสังคม (Social Criteria)	0.081	
- สุขภาพและความปลอดภัย (Health and Safety)		0.750
- การรับผิดชอบต่อสังคม (Corporate Social Responsibility)		0.250

ที่มา: (Chaitongrat & Areerakulkan, 2021)

ตารางที่ 12.2 คะแนนประเมินซัพพลายเออร์ทั้ง 4 ราย

ผู้ขาย	ค่าน้ำหนักทางเลือก
ผู้ขาย A	0.155
ผู้ขาย B	0.357
ผู้ขาย C	0.232
ผู้ขาย D	0.256

ที่มา: (Chaitongrat & Areerakulkan, 2021)

ผลการวิเคราะห์พบว่า ค่าคะแนนความสอดคล้องเฉลี่ยมีค่าเท่ากับ 0.07 แสดงว่าผู้เชี่ยวชาญได้ทำการเปรียบเทียบรายคู่ของเกณฑ์ต่างๆและให้คะแนนประเมินสอดคล้องกัน จากตารางที่ 12.2 ซัพพลายเออร์ที่มีค่าคะแนนสูงสุดคือ ซัพพลายเออร์ B มีค่าคะแนนเท่ากับ 0.357 รองลงมาคือ ซัพพลายเออร์ D มีค่าคะแนนเท่ากับ 0.256 ซัพพลายเออร์ C มีค่าคะแนนเท่ากับ 0.232 สุดท้ายคือ ซัพพลายเออร์ A ที่มีค่าคะแนนเท่ากับ 0.155 จึงสรุปได้ว่า ซัพพลายเออร์ที่เหมาะสมกับเกณฑ์การประเมินซัพพลายเออร์อย่างยั่งยืนคือ ซัพพลายเออร์ B เนื่องจากมีค่าน้ำหนักประเมินสูงที่สุด

โดยสรุปกรณีศึกษาการคัดเลือกผู้ขายอย่างยั่งยืนด้วยวิธีการวิเคราะห์เชิงลำดับชั้น (AHP) นี้แสดงให้เห็นถึงการประยุกต์ใช้เครื่องมือเชิงวิเคราะห์ในการแก้ปัญหาการตัดสินใจที่มีความซับซ้อนในบริบทของการจัดการซัพพลายเชน โดยเฉพาะปัญหาที่เกี่ยวข้องกับการพิจารณาหลายเกณฑ์พร้อมกันทั้งด้านเศรษฐกิจ สิ่งแวดล้อม และสังคม ซึ่งเป็นองค์ประกอบสำคัญของแนวคิดความยั่งยืน กรณีศึกษานี้ช่วยให้เข้าใจถึงกระบวนการตัดสินใจอย่างเป็นระบบ ตั้งแต่การกำหนดปัญหา การสร้างโครงสร้างลำดับชั้น การรวบรวมข้อมูลจากผู้เชี่ยวชาญ การเปรียบเทียบเชิงคู่ การคำนวณค่าน้ำหนัก และการจัดอันดับทางเลือก

การนำกรณีศึกษานี้ไปใช้งานสามารถช่วยพัฒนาทักษะด้านการคิดเชิงวิเคราะห์และการตัดสินใจเชิงข้อมูล (data-driven decision making) โดยสามารถประยุกต์ใช้กับปัญหาทางธุรกิจอื่นๆ ที่มีลักษณะคล้ายคลึงกัน เช่น การคัดเลือกซัพพลายเออร์ในอุตสาหกรรมต่าง ๆ การเลือกทำเลที่ตั้ง การประเมินความเสี่ยง หรือการจัดลำดับความสำคัญของโครงการ นอกจากนี้ แนวทางดังกล่าวยังสามารถต่อยอดไปสู่การใช้เทคนิคขั้นสูงร่วมกับ AHP เช่น การรวมกับแบบจำลองฟัซซี (Fuzzy AHP) หรือวิธีการจัดอันดับอื่น ๆ เพื่อรองรับความไม่แน่นอนและเพิ่มความแม่นยำของผลการวิเคราะห์ (Sarkis & Dhavale, 2015) ท้ายที่สุดกรณีศึกษานี้เป็นตัวอย่างที่สะท้อนให้เห็นถึงความสำคัญของการใช้เครื่องมือเชิงวิเคราะห์ในการสนับสนุนการตัดสินใจในซัพพลายเชน ซึ่งไม่เพียงช่วยเพิ่มประสิทธิภาพในการดำเนินงาน แต่ยังช่วยให้องค์กรสามารถดำเนินธุรกิจภายใต้กรอบแนวคิดความยั่งยืนได้อย่างเป็นรูปธรรม และสามารถนำไปประยุกต์ใช้ในสถานการณ์จริงได้อย่างมีประสิทธิภาพและเป็นระบบ

เอกสารอ้างอิง

- Amindoust, A., Ahmed, S., Saghafinia, A., & Bahreininejad, A. (2012). Sustainable supplier selection: A ranking model based on fuzzy inference system. *Applied Soft Computing*, 12(6), 1668-1677. <https://doi.org/10.1016/j.asoc.2012.01.023>.
- Chaitongrat, S., & Areerakulkan, N. (2021). Application of analytic hierarchy process method for sustainable supplier selection: A case study of ABC Company. *Journal of Logistics and Supply Chain College*, 7(2), 5-17.
- Holden, E., Linnerud, K., & Banister, D. (2014). Sustainable development: Our common future revisited. *Global Environmental Change*, 26, 130-139. <https://doi.org/10.1016/j.gloenvcha.2014.04.006>.
- Memari, A., Dargi, A., Jokar, M. R. A., Ahmad, R., & Rahim, A. R. A. (2019). Sustainable supplier selection: A multi-criteria intuitionistic fuzzy TOPSIS method. *Journal of Manufacturing Systems*, 50, 9-24. <https://doi.org/10.1016/j.jmsy.2018.11.002>.
- Pishchulov, G., Trautrim, A., Chesney, T., Gold, S., & Schwab, L. (2019). The voting analytic hierarchy process revisited: A revised method with application to sustainable supplier selection. *International Journal of Production Economics*, 211, 166-179. <https://doi.org/10.1016/j.ijpe.2019.01.025>.
- Purvis, B., Mao, Y., & Robinson, D. (2019). Three pillars of sustainability: In search of conceptual origins. *Sustainability Science*, 14, 681-695. <https://doi.org/10.1007/s11625-018-0627-5>.
- Saaty, T. L. (1980). *The analytic hierarchy process*. McGraw-Hill.
- Sarkis, J., & Dhavale, D.G. (2015). Supplier selection for sustainable operations: A triple-bottom-line approach using a Bayesian framework. *International Journal of Production Economics*, 166, 177-191. <https://doi.org/10.1016/j.ijpe.2014.11.007>

บทที่ 13

งานวิจัยที่น่าสนใจเกี่ยวกับการพยากรณ์สินค้าใหม่ด้วย วิธีโครงข่ายประสาทเทียม

Interesting Research on New Product Demand Forecasting Using Artificial Neural Networks

หัวข้อ

- 13.1 บริบทและปัญหาของกรณีศึกษา
- 13.2 การวิเคราะห์ข้อมูลและพัฒนาแบบจำลอง
- 13.3 การประเมินและเปรียบเทียบผลการพยากรณ์
- 13.4 การประยุกต์ใช้ผลการพยากรณ์ในซัพพลายเชน
- 13.5 บทเรียนจากงานวิจัยและสรุป

การพยากรณ์อุปสงค์เป็นกระบวนการสำคัญในการบริหารจัดการซัพพลายเชน เนื่องจากช่วยให้องค์กรสามารถวางแผนการจัดซื้อ การผลิต การจัดการสินค้าคงคลัง และการกระจายสินค้าได้อย่างมีประสิทธิภาพ ความแม่นยำของการพยากรณ์มีผลโดยตรงต่อต้นทุนการดำเนินงาน ระดับการให้บริการลูกค้า และความสามารถในการแข่งขันขององค์กร โดยเฉพาะอย่างยิ่งในกรณีของการเปิดตัวผลิตภัณฑ์ใหม่ ซึ่งมักมีความไม่แน่นอนและมีข้อมูลในอดีตสำหรับการพยากรณ์อย่างจำกัด

ในทางปฏิบัติ การพยากรณ์อุปสงค์ที่คลาดเคลื่อนอาจนำไปสู่ปัญหาการผลิตสินค้าไม่เพียงพอต่อความต้องการ การขาดแคลนสินค้า การสูญเสียโอกาสทางการขาย หรือการมีสินค้าคงคลังมากเกินไป ความจำเป็น ส่งผลให้ต้นทุนการดำเนินงานเพิ่มสูงขึ้นและประสิทธิภาพของซัพพลายเชนลดลง ดังนั้น การเลือกใช้วิธีการพยากรณ์ที่เหมาะสมจึงเป็นปัจจัยสำคัญที่ช่วยสนับสนุนการตัดสินใจเชิงกลยุทธ์และเชิงปฏิบัติการขององค์กร

ปัจจุบัน การประยุกต์ใช้การเรียนรู้ของเครื่อง (ML) ได้รับความสนใจอย่างแพร่หลายในการพยากรณ์อุปสงค์ เนื่องจากสามารถเรียนรู้รูปแบบความสัมพันธ์ที่ซับซ้อนของข้อมูลและปรับตัวต่อการเปลี่ยนแปลงของสภาพแวดล้อมทางธุรกิจได้ดีกว่าวิธีการทางสถิติแบบดั้งเดิม หนึ่งในเทคนิคที่ได้รับ

ความนิยมคือ โครงข่ายประสาทเทียม (ANN) ซึ่งสามารถนำมาใช้ในการสร้างแบบจำลองการพยากรณ์ที่มีความยืดหยุ่นและให้ผลลัพธ์ที่มีความแม่นยำสูง (Haykin, 2009; Vandepu, 2021).)

บทนี้นำเสนอกรณีศึกษาการพยากรณ์อุปสงค์ผลิตภัณฑ์ใหม่ด้วยโครงข่ายประสาทเทียม อ้างอิงจากการศึกษาของ Areerakulkan et al. (2024) โดยเริ่มจากการวิเคราะห์ปัญหาทางธุรกิจ และข้อมูลที่เกี่ยวข้อง จากนั้นพัฒนาและประเมินแบบจำลองการพยากรณ์ พร้อมเปรียบเทียบผลลัพธ์กับวิธีการพยากรณ์แบบดั้งเดิม ตลอดจนแสดงการประยุกต์ใช้ผลการพยากรณ์ในการวางแผนการผลิต และการจัดการสินค้าคงคลัง เพื่อสนับสนุนการตัดสินใจและเพิ่มประสิทธิภาพของซัพพลายเชน อันเป็นการเชื่อมโยงองค์ความรู้ด้านการวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเข้ากับการประยุกต์ใช้ในสถานการณ์จริงทางธุรกิจ

13.1 บริบทและปัญหาของกรณีศึกษา

อุตสาหกรรมโทรศัพท์เคลื่อนที่เป็นหนึ่งในอุตสาหกรรมที่มีการแข่งขันสูงและมีการเปลี่ยนแปลงของเทคโนโลยีอย่างรวดเร็ว ผู้ผลิตต้องเปิดตัวผลิตภัณฑ์ใหม่อย่างต่อเนื่องเพื่อตอบสนองต่อความต้องการของผู้บริโภคและรักษาความสามารถในการแข่งขัน ส่งผลให้ผลิตภัณฑ์มีวงจรชีวิตสั้น (short product life cycle) และมีความไม่แน่นอนของอุปสงค์สูง โดยยอดขายมักเพิ่มขึ้นอย่างรวดเร็วในช่วงเปิดตัวผลิตภัณฑ์ ก่อนจะค่อย ๆ ลดลงเมื่อมีผลิตภัณฑ์รุ่นใหม่เข้าสู่ตลาด (Christopher, 2016)

นอกจากข้อจำกัดด้านวงจรชีวิตของผลิตภัณฑ์แล้ว การพยากรณ์อุปสงค์ของโทรศัพท์เคลื่อนที่ก็ยังได้รับผลกระทบจากหลายปัจจัย เช่น การเปลี่ยนแปลงของเทคโนโลยี การแข่งขันด้านราคา กิจกรรมส่งเสริมการขาย การเปิดตัวผลิตภัณฑ์ของคู่แข่ง และพฤติกรรมของผู้บริโภคที่เปลี่ยนแปลงอย่างรวดเร็ว ปัจจัยเหล่านี้ทำให้รูปแบบของอุปสงค์มีทั้งแนวโน้ม (trend) และความผันผวนตามช่วงเวลา (seasonality) ส่งผลให้การวางแผนการผลิตและการจัดการสินค้าคงคลังมีความซับซ้อนมากขึ้น (Ingle et al., 2021)

กรณีศึกษาอ้างอิงจากการศึกษาของ Areerakulkan et al. (2024) ซึ่งเป็นบริษัทผู้ผลิตโทรศัพท์เคลื่อนที่ที่ประสบปัญหาในการวางแผนอุปทานสำหรับการเปิดตัวผลิตภัณฑ์รุ่นใหม่ โดยบริษัทใช้การพยากรณ์อุปสงค์เป็นข้อมูลพื้นฐานในการกำหนดแผนการผลิต การเตรียมสินค้าสำเร็จรูป การผลิตสินค้าระหว่างกระบวนการ และการจัดสรรกำลังการผลิต อย่างไรก็ตาม ความคลาดเคลื่อนของการพยากรณ์ส่งผลให้ปริมาณสินค้าที่เตรียมไว้ไม่สอดคล้องกับความต้องการของตลาด

ปัญหาที่สำคัญประการแรก คือ การขาดแคลนสินค้าในช่วงเปิดตัวผลิตภัณฑ์ เนื่องจากปริมาณการผลิตต่ำกว่าความต้องการที่เกิดขึ้นจริง ส่งผลให้ลูกค้าต้องรอรับสินค้าเป็นเวลานาน เกิด

การยกเลิกคำสั่งซื้อ และสูญเสียโอกาสทางการขายให้กับคู่แข่ง ซึ่งส่งผลกระทบต่อรายได้และภาพลักษณ์ขององค์กร

ประการที่สอง ความคลาดเคลื่อนของการพยากรณ์ส่งผลต่อการวางแผนการผลิตและการจัดการสินค้าคงคลังตลอดทั้งซัพพลายเชน เนื่องจากการกำหนดปริมาณสินค้าสำเร็จรูป (finished goods, FG) และสินค้าระหว่างผลิต (work in process, WIP) อาศัยข้อมูลการพยากรณ์เป็นหลัก เมื่อค่าพยากรณ์ต่ำกว่าความเป็นจริง องค์กรไม่สามารถจัดเตรียมสินค้าคงคลังและกำลังการผลิตได้เพียงพอต่อความต้องการของตลาด ส่งผลให้เกิดปัญหาการขาดแคลนสินค้าต่อเนื่องในหลายช่วงเวลา

ประการที่สาม บริษัทต้องเผชิญกับข้อจำกัดด้านกำลังการผลิต (capacity constraint) เนื่องจากการเพิ่มกำลังการผลิตในระยะสั้นสามารถทำได้จำกัด การตัดสินใจเกี่ยวกับการผลิตจึงต้องอาศัยการพยากรณ์อุปสงค์ที่มีความแม่นยำ หากการคาดการณ์ผิดพลาด องค์กรอาจไม่สามารถตอบสนองต่อความต้องการของตลาดได้ทันเวลา หรือในทางกลับกัน อาจเกิดกำลังการผลิตและสินค้าคงคลังส่วนเกิน ซึ่งเป็นภาระด้านต้นทุน

ประการที่สี่ เนื่องจากโทรศัพท์เคลื่อนที่มีวงจรชีวิตสั้น สินค้าคงคลังที่เหลือจากการวางแผนที่ผิดพลาดมีความเสี่ยงต่อการล้าสมัยและการลดมูลค่าของสินค้าอย่างรวดเร็ว โดยเฉพาะเมื่อมีการเปิดตัวผลิตภัณฑ์รุ่นใหม่เข้าสู่ตลาด องค์กรอาจจำเป็นต้องลดราคาเพื่อระบายสินค้า ส่งผลกระทบต่อกำไรและผลตอบแทนจากการลงทุน (Simchi-Levi et al., 2021)

จากการวิเคราะห์ปัญหาของกรณีศึกษา พบว่าปัญหาหลักไม่ได้เกิดจากการขาดแคลนกำลังการผลิตหรือสินค้าคงคลังเพียงอย่างเดียว แต่มีสาเหตุสำคัญมาจากความคลาดเคลื่อนของการพยากรณ์อุปสงค์ ซึ่งส่งผลกระทบเป็นลูกโซ่ต่อการวางแผนการผลิต การจัดการสินค้าคงคลัง และการบริหารทรัพยากรในซัพพลายเชน ดังนั้น การพัฒนาแบบจำลองการพยากรณ์ที่มีความแม่นยำสูงจึงเป็นแนวทางสำคัญในการลดความไม่แน่นอนของการดำเนินงาน เพิ่มระดับการให้บริการลูกค้า และลดต้นทุนรวมของซัพพลายเชน

เพื่อตอบเจตน์ดังกล่าว งานวิจัยนี้จึงมุ่งศึกษาการประยุกต์ใช้โครงข่ายประสาทเทียม (ANN) และเปรียบเทียบประสิทธิภาพกับวิธีการพยากรณ์แบบดั้งเดิม เพื่อค้นหาแบบจำลองที่เหมาะสมสำหรับการพยากรณ์อุปสงค์ของผลิตภัณฑ์ใหม่ และนำผลการพยากรณ์ไปใช้สนับสนุนการวางแผนการผลิต การจัดการสินค้าคงคลัง และการตัดสินใจด้านซัพพลายเชนอย่างมีประสิทธิภาพ

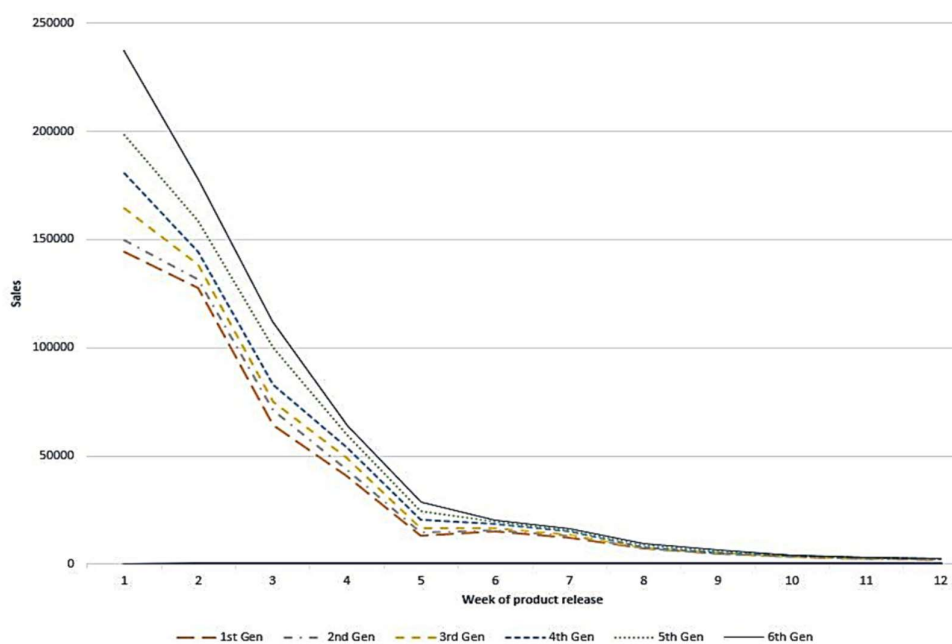
13.2 การวิเคราะห์ข้อมูลและพัฒนาแบบจำลอง

หลังจากศึกษาบริบทและปัญหาของกรณีศึกษาแล้ว ขั้นตอนสำคัญถัดไปคือการวิเคราะห์ข้อมูลและการพัฒนาแบบจำลองการพยากรณ์อุปสงค์ เพื่อทำความเข้าใจลักษณะของข้อมูลและเลือก

วิธีการพยากรณ์ที่เหมาะสมสำหรับผลิตภัณฑ์ใหม่ที่มีความไม่แน่นอนของอุปสงค์สูง โดยกระบวนการดังกล่าวประกอบด้วยการสำรวจข้อมูล การเตรียมข้อมูล และการพัฒนาแบบจำลองการพยากรณ์

13.2.1 การสำรวจและวิเคราะห์ข้อมูล

กรณีศึกษาที่ใช้ข้อมูลยอดขายรายสัปดาห์ของโทรศัพท์เคลื่อนที่จากผลิตภัณฑ์รุ่นก่อนหน้าจำนวน 6 รุ่นผลิตภัณฑ์ โดยใช้ข้อมูลของรุ่นที่ 1–5 สำหรับการสร้างแบบจำลอง และใช้ข้อมูลของรุ่นที่ 6 สำหรับการทดสอบประสิทธิภาพของแบบจำลอง (Areerakulkan et al., 2024) ก่อนการพัฒนาแบบจำลอง มีการสำรวจข้อมูลเบื้องต้น (exploratory data analysis, EDA) เพื่อศึกษาลักษณะของอุปสงค์ พบว่าอุปสงค์ของโทรศัพท์เคลื่อนที่มีรูปแบบเฉพาะของผลิตภัณฑ์ที่มีวงจรชีวิตสั้น กล่าวคือ ยอดขายจะเพิ่มสูงมากในช่วงเริ่มต้นของการเปิดตัวสินค้า จากนั้นจะลดลงอย่างรวดเร็วในช่วงแรก และค่อย ๆ ลดลงจนเข้าสู่ระดับคงที่ในระยะต่อมา รูปแบบดังกล่าวสะท้อนถึงองค์ประกอบของแนวโน้ม (trend) และฤดูกาล (seasonality) ที่ควรนำมาพิจารณาในการเลือกแบบจำลองการพยากรณ์



ภาพที่ 13.1 กราฟข้อมูลยอดขายรายสัปดาห์ของโทรศัพท์เคลื่อนที่จากผลิตภัณฑ์รุ่นก่อนหน้า

จำนวน 6 รุ่นผลิตภัณฑ์

ที่มา: (Areerakulkan et al., 2021)

จากนั้นจึงทำการเปรียบเทียบค่าพยากรณ์เดิมของบริษัทกับยอดขายจริง (รุ่นที่ 6) เพื่อประเมินปัญหาที่เกิดขึ้น พบว่าค่าพยากรณ์ต่ำกว่าความต้องการจริงในหลายช่วงเวลา ส่งผลให้เกิดการขาดแคลนสินค้าและความผิดพลาดในการวางแผนการผลิตและสินค้าคงคลัง ยอดพยากรณ์แสดงดังตารางต่อไปนี้

ตารางที่ 13.1 เปรียบเทียบค่าพยากรณ์และค่าจริง

Week	Demand (Units)	Forecast (Units)	Deviation
50	237,450	210,000	-27,450
51	177,440	150,000	-27,440
52	112,116	100,000	-12,116
1	63,883	75,000	11,117
2	28,614	26,000	-2,614
3	20,573	15,000	-5,573
4	16,408	11,000	-5,408
5	9,550	8,500	-1,050
6	6,561	5,000	-1,561
7	4,159	3,500	-659
8	3,108	2,600	-508
9	2,770	2,600	-170

หมายเหตุ: Demand คือ ยอดอุปสงค์จริงในอดีตของ Gen 6th Forecast คือ ยอดพยากรณ์ด้วยวิธีเดิมของ Gen 6th และ Deviation คือ ค่าเบี่ยงเบนจากยอดพยากรณ์เมื่อเปรียบเทียบกับค่าจริง ตัวเลขติดลบหมายถึง พยากรณ์น้อยกว่าที่ควรจะเป็น

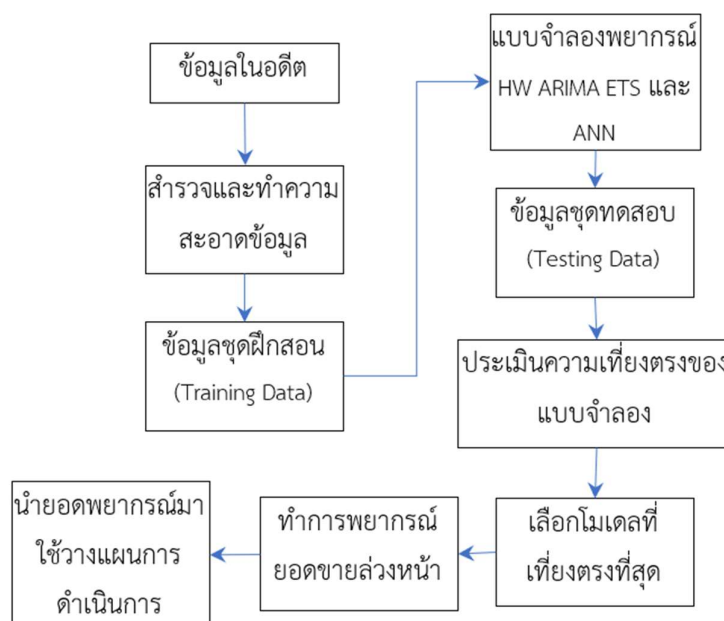
ที่มา: (Areerakulkan et al., 2021)

13.2.2 การพัฒนาแบบจำลองการพยากรณ์

จากการทบทวนวรรณกรรมและลักษณะของข้อมูลพบว่าวิธีการพยากรณ์ที่เหมาะสมสำหรับการศึกษานี้สามารถแบ่งออกเป็นสองกลุ่ม ได้แก่ วิธีการทางสถิติแบบดั้งเดิม และวิธีการเรียนรู้ของเครื่อง (Ingle et al., 2021) เพื่อเปรียบเทียบประสิทธิภาพของแบบจำลอง จึงเลือกวิธีการพยากรณ์จำนวน 4 วิธี ได้แก่ Holt-Winters (HW) Autoregressive Integrated Moving Average (ARIMA) Exponential Smoothing (ETS) และ Artificial Neural Network (ANN) แบบจำลองทั้งสี่ถูกพัฒนาขึ้นโดยใช้ชุดข้อมูลเดียวกัน เพื่อให้สามารถเปรียบเทียบประสิทธิภาพได้อย่างเป็นธรรม โดยใช้ข้อมูลในอดีตสำหรับการเรียนรู้รูปแบบของอุปสงค์ และสร้างค่าพยากรณ์สำหรับข้อมูลชุดทดสอบ

สำหรับโครงข่ายประสาทเทียม (ANN) ซึ่งเป็นการประยุกต์ใช้โครงข่ายประสาทเทียมกับข้อมูลอนุกรมเวลา โดยใช้ข้อมูลยอดขายย้อนหลังเป็นข้อมูลนำเข้าเพื่อเรียนรู้รูปแบบของอุปสงค์และสร้างค่าพยากรณ์ในอนาคต

กระบวนการพัฒนาแบบจำลองเริ่มจากการกำหนดข้อมูลนำเข้า การฝึกแบบจำลองด้วยชุดข้อมูลฝึก และการปรับค่าพารามิเตอร์เพื่อให้ความคลาดเคลื่อนของการพยากรณ์มีค่าน้อยที่สุด จากนั้นจึงนำแบบจำลองที่ได้ไปทดสอบกับข้อมูลชุดทดสอบ เพื่อประเมินความสามารถในการพยากรณ์ของแต่ละวิธี และเลือกใช้โมเดลที่เที่ยงตรงมากที่สุดในการพยากรณ์อุปสงค์ล่วงหน้าของมือถือ generation ถัดไป หลังจากนั้นนำยอดพยากรณ์ไปใช้ในการวางแผนการดำเนินงาน ขั้นตอนของการพัฒนาแบบจำลองการพยากรณ์แสดงดังภาพที่ 13.2



ภาพที่ 13.2 ขั้นตอนการพัฒนาโมเดลพยากรณ์

ที่มา: (Areerakulkan et al., 2021)

การพัฒนาแบบจำลองทั้งสี่แบบในกรณีศึกษานี้มีวัตถุประสงค์เพื่อเปรียบเทียบความสามารถในการพยากรณ์อุปสงค์ของผลิตภัณฑ์ใหม่ และคัดเลือกแบบจำลองที่มีความเหมาะสมที่สุดสำหรับการประยุกต์ใช้ในการวางแผนการผลิตและการจัดการสินค้าคงคลังในซัพพลายเชน ซึ่งผลการประเมินประสิทธิภาพของแบบจำลองจะนำเสนอในหัวข้อถัดไป

13.3 การประเมินและเปรียบเทียบผลการพยากรณ์

หลังจากการพัฒนาแบบจำลองการพยากรณ์ ขั้นตอนสำคัญถัดไปคือการประเมินประสิทธิภาพและเปรียบเทียบผลการพยากรณ์ของแต่ละแบบจำลอง เพื่อคัดเลือกวิธีการที่มีความเหมาะสมสำหรับการพยากรณ์อุปสงค์ของผลิตภัณฑ์ใหม่ โดยการประเมินประสิทธิภาพของแบบจำลองเป็นขั้นตอนที่ช่วยตรวจสอบว่าแบบจำลองสามารถอธิบายรูปแบบของข้อมูลและพยากรณ์อุปสงค์ในอนาคตได้อย่างแม่นยำเพียงใด

13.3.1 เกณฑ์การประเมินแบบจำลอง

ในการศึกษานี้ ใช้ตัวชี้วัดความคลาดเคลื่อนของการพยากรณ์ที่ได้รับความนิยม ได้แก่ mean absolute deviation (MAD) เป็นค่าเฉลี่ยของความคลาดเคลื่อนสัมบูรณ์ระหว่างค่าพยากรณ์และค่าจริง โดยค่าที่ต่ำกว่าหมายถึงแบบจำลองมีความแม่นยำสูงกว่า root mean square error (RMSE) เป็นค่ารากที่สองของค่าเฉลี่ยกำลังสองของความคลาดเคลื่อน ซึ่งให้ความสำคัญกับความผิดพลาดขนาดใหญ่เป็นพิเศษ mean absolute percentage error (MAPE) เป็นค่าร้อยละของความคลาดเคลื่อนสัมบูรณ์เฉลี่ย ซึ่งช่วยให้สามารถเปรียบเทียบประสิทธิภาพของแบบจำลองได้ง่าย (Areerakulkan et al., 2024)

13.3.2 การเปรียบเทียบผลการพยากรณ์

เพื่อประเมินประสิทธิภาพของแบบจำลองในเชิงปฏิบัติ ได้ทำการเปรียบเทียบค่าพยากรณ์ของแบบจำลองทั้ง 4 แบบ กับวิธีการพยากรณ์เดิมของบริษัท โดยใช้ข้อมูลของผลิตภัณฑ์รุ่นที่ 1-5 ผลของการเปรียบเทียบแสดงดังตารางที่ 13.2

ตารางที่ 13.2 เปรียบเทียบค่าพยากรณ์ของแบบจำลองทั้ง 4 แบบ

ตัวชี้วัด ความคลาดเคลื่อน	แบบจำลอง			
	HW	ARIMA	ETS	ANN
MAD	1,202.20	1,280.54	1,210.91	949.97
RMSE	2,147.60	2,180.47	2,354.12	1,427.46
MAPE	4.89	7.84	2.99	5.41

ที่มา: (Areerakulkan et al., 2021)

ผลการศึกษาแสดงให้เห็นว่า ค่าพยากรณ์จาก ANN มีความใกล้เคียงกับยอดขายจริงมากกว่าวิธีการเดิม โดยเฉพาะในช่วงที่มีความต้องการสูงในระยะเริ่มต้นของการเปิดตัวผลิตภัณฑ์ ซึ่งเป็นช่วงเวลาที่มีความสำคัญต่อการวางแผนการผลิตและการจัดการสินค้าคงคลัง จากตารางที่ 13.2 ยังพบว่า แบบจำลอง ANN ให้ค่าความคลาดเคลื่อนที่น้อยที่สุดสำหรับ MAD RMSE แต่สำหรับ MAPE แบบจำลอง ETS จะให้ค่าที่น้อยที่สุด สำหรับงานวิจัยนี้จะใช้ RMSE เป็นหลักเนื่องจาก RMSE จะให้ค่าพยากรณ์ประมาณค่าเฉลี่ยแทนที่จะมากเกินไปหรือน้อยเกินไป ด้วยเหตุนี้เองแบบจำลอง ANN จึงเป็นตัวเลือกที่ดีที่สุดสำหรับพยากรณ์ข้อมูลชุดทดสอบ (testing set, 6th generation) และเปรียบเทียบกับวิธีการพยากรณ์แบบเดิม ผลของการพยากรณ์แสดงดังตารางที่ 13.3

ตารางที่ 13.3 เปรียบเทียบค่าพยากรณ์ของมือถือรุ่นที่ 6 ระหว่างวิธี ANN กับวิธีเดิม

Week	Demand	Old Forecast	New Forecast (ANN.)
1	237,450	210,000	219,496
2	177,440	150,000	174,983
3	112,116	100,000	122,044
4	63,883	75,000	67,042
5	28,614	26,000	28,560
6	20,573	15,000	21,117
7	16,408	11,000	15,413
8	9,550	8,500	8,988
9	6,561	5,000	6,431
10	4,159	3,500	4,325
11	3,108	2,600	3,322
12	2,770	2,600	2,968

หมายเหตุ: Demand คือ ยอดอุปสงค์จริงในอดีตของ Gen 6th Old Forecast คือ ยอดพยากรณ์ด้วยวิธีเดิมของ Gen 6th และ New Forecast (ANN.) คือ ยอดพยากรณ์ด้วยแบบจำลอง ANN.

ที่มา: (Areerakulkan et al., 2021)

จากตารางที่ 13.3 พบว่า วิธีการพยากรณ์แบบเดิมมีค่าความคลาดเคลื่อน MAD = 7,972.17 RMSE = 12,410.48 และ MAPE = 16.46 ในขณะที่แบบจำลองการพยากรณ์ ANN มีค่าความคลาดเคลื่อน MAD = 3,030.04 RMSE = 6,046.08 และ MAPE = 5.41 หรืออาจกล่าวได้ว่าแบบจำลอง ANN 51.28% เที่ยงตรงมากกว่าการพยากรณ์ด้วยวิธีเดิม

13.4 การประยุกต์ใช้ผลการพยากรณ์ในซัพพลายเชน

การพยากรณ์อุปสงค์ไม่ได้มีวัตถุประสงค์เพียงเพื่อคาดการณ์ยอดขายในอนาคตเท่านั้น แต่ยังเป็นข้อมูลพื้นฐานที่สำคัญสำหรับการตัดสินใจด้านการบริหารซัพพลายเชน ทั้งในระดับกลยุทธ์ ยุทธวิธี และการปฏิบัติการ โดยความแม่นยำของการพยากรณ์มีผลโดยตรงต่อการวางแผนการผลิต การจัดการสินค้าคงคลัง การจัดซื้อวัตถุดิบ การจัดสรรกำลังการผลิต และการกระจายสินค้า (Simchi-Levi et al., 2021)

จากผลการศึกษาในหัวข้อก่อนหน้านี้ พบว่าแบบจำลอง ANN มีประสิทธิภาพในการพยากรณ์อุปสงค์ของผลิตภัณฑ์ใหม่ได้ดีกว่าวิธีการพยากรณ์แบบเดิม จึงสามารถนำผลการพยากรณ์ไปประยุกต์ใช้เพื่อสนับสนุนการตัดสินใจในกิจกรรมต่าง ๆ ของซัพพลายเชนเริ่มจาก การวางแผนการผลิตและการจัดการสินค้าคงคลัง

13.4.1 การวางแผนการผลิตและการจัดการสินค้าคงคลัง

สำหรับการวางแผนการผลิตเริ่มจากการคำนวณจำนวนสัปดาห์ในการสต็อกสินค้าสำเร็จรูป (finish goods, FG) รวมถึงสินค้าในกระบวนการ (work in process, WIP) โดยใช้สมการที่ 13.1

$$n = \frac{\sum_{w=1}^{k+1} F_w}{Cap_{fg/wip}} \quad 13.1$$

โดยที่ F_w = ค่าพยากรณ์ของอุปสงค์ที่เกิน capacity

w = สัปดาห์ที่ค่าพยากรณ์เกิน capacity

k = จำนวนสัปดาห์ที่อุปสงค์เกิน capacity

n = จำนวนสัปดาห์ที่จำเป็นในการสต็อก FG

จากการคำนวณด้วยสมการที่ 13.1 พบว่า WIP ควรเริ่มผลิตตั้งแต่สัปดาห์ที่ 37 และ FG ควรเริ่มผลิตตั้งแต่สัปดาห์ที่ 42 ตามลำดับ นอกจากนี้ capacity ควรเพิ่มจาก 42,000 หน่วยเป็น 50,000 หน่วย โดยแผนการผลิตแสดงดังตารางที่ 13.4

ตารางที่ 13.4 การวางแผนการผลิตเริ่มตั้งแต่สัปดาห์ที่ 37 ไปจนถึงสัปดาห์ที่ 9 ของปีปัจจุบัน

Events	Week	Pr. Quantity (FG.)	Stocks (FG.)	Pr. Quantity (WIP.)	Stocks (WIP.)
Start WIP.	37	0	0	37,500	37,500
	38	0	0	37,500	75,000
	39	0	0	37,500	112,500
	40	0	0	37,500	150,000
	41	0	0	37,500	187,500
Start FG.	42	50,000	50,000	37,500	175,000
	43	50,000	100,000	37,500	162,500
	44	50,000	150,000	37,500	150,000
	45	50,000	200,000	37,500	137,500
	46	50,000	250,000	37,500	125,000
	47	50,000	300,000	37,500	112,500
	48	50,000	350,000	37,500	100,000
	49	50,000	400,000	37,500	87,500
	50	50,000	215,000	37,500	75,000
Release FG.	51	50,000	90,000	37,500	62,500
	52	50,000	30,000	37,500	50,000
	1	50,000	18,000	37,500	37,500
	2	27,061	18,061	0	10,439
	3	10,439	9,500	18,944	18,944
	4	14,974	9,474	0	3,970
	5	3,970	4,944	8,237	8,237
	6	5,409	4,853	0	2,828
	7	2,828	4,181	3,423	3,423
	8	2,355	4,086	0	1,068
	9	1,068	2,604	2,462	2,462

หมายเหตุ: Pr. Quantity (FG.) คือ ปริมาณสั่งผลิตสินค้าสำเร็จรูป Stocks (FG.) คือ สต็อกสินค้าสำเร็จรูป Pr. Quantity (WIP.) คือ ปริมาณสินค้าระหว่างกระบวนการผลิต และ Stocks (WIP.) สต็อกสินค้าระหว่างกระบวนการผลิต

ที่มา: (Areerakulkan et al., 2021)

จากการวางแผนการผลิตในตารางที่ 13.4 ถ้าทำการประมาณต้นทุนบนพื้นฐานของส่วนประกอบต้นทุน 2 ชนิดได้แก่ ต้นทุนสินค้าขาดและต้นทุนเก็บรักษาสินค้าคงคลัง พบว่าถ้ากำหนดให้กำไร (sales margin) สำหรับกรณีนี้เท่ากับ (\$155.00) ต้นทุนสินค้าขาดสำหรับปัญหา ก่อนปรับปรุงเท่ากับ \$3,085,070.00 ในขณะที่ ต้นทุนสินค้าขาดหลังปรับปรุงเป็น \$0.00 หรือกล่าวอีกนัยหนึ่งคือลูกค้าได้รับสินค้าครบ 100% สำหรับต้นทุนเก็บรักษาสต็อกคำนวณโดยใช้ ต้นทุนเก็บรักษา สำหรับ $WIP = \$0.53$ และ ต้นทุนเก็บรักษาสำหรับ $FG = \$1.20$ จะสามารถคำนวณต้นทุนเก็บรักษาทั้งหมดก่อนปรับปรุง = \$1,969,441.60 ในขณะที่ ต้นทุนเก็บรักษาทั้งหมดหลังปรับปรุง = \$3,653,885.23 ถ้ารวมต้นทุนทั้ง 2 ชนิดเข้าด้วยกันจะพบว่าต้นทุนรวม ก่อนปรับปรุง = \$5,054,512.03 และ ต้นทุนรวมหลังปรับปรุง = \$1,400,626.80 หรือลดลง 27.71%

13.5 บทเรียนจากกรณีศึกษาและสรุป

งานวิจัยนี้แสดงให้เห็นถึงความสำคัญของการพยากรณ์อุปสงค์ต่อประสิทธิภาพของการจัดการซัพพลายเชน โดยเฉพาะในอุตสาหกรรมโทรศัพท์เคลื่อนที่ที่มีการแข่งขันสูง มีการเปลี่ยนแปลงทางเทคโนโลยีอย่างรวดเร็ว และมีวงจรชีวิตผลิตภัณฑ์สั้น ความคลาดเคลื่อนของการพยากรณ์อุปสงค์ส่งผลกระทบโดยตรงต่อการวางแผนการผลิต การจัดการสินค้าคงคลัง และระดับการให้บริการลูกค้า ซึ่งอาจนำไปสู่ปัญหาสินค้าขาดแคลน การสูญเสียยอดขาย และต้นทุนการดำเนินงานที่เพิ่มขึ้น

จากการศึกษาพบว่า ปัญหาหลักของบริษัทไม่ได้เกิดจากข้อจำกัดด้านกำลังการผลิตหรือการจัดการสินค้าคงคลังเพียงอย่างเดียว แต่มีสาเหตุสำคัญมาจากความคลาดเคลื่อนของการพยากรณ์อุปสงค์ที่ใช้ในการวางแผน เมื่อค่าพยากรณ์ต่ำกว่าความต้องการที่เกิดขึ้นจริง บริษัทไม่สามารถเตรียมสินค้าสำเร็จรูปและสินค้าระหว่างผลิตได้เพียงพอต่อความต้องการของตลาด ส่งผลให้เกิดการขาดแคลนสินค้าและการสูญเสียโอกาสทางการขาย

ผลการเปรียบเทียบแบบจำลองการพยากรณ์แสดงให้เห็นว่า โครงข่ายประสาทเทียม (ANN) มีประสิทธิภาพสูงกว่าวิธีการพยากรณ์แบบดั้งเดิม ได้แก่ Holt-Winters (HW) ARIMA และ Exponential Smoothing (ETS) โดยสามารถลดค่าความคลาดเคลื่อนของการพยากรณ์และให้ผลลัพธ์ที่ใกล้เคียงกับความต้องการจริงมากกว่า ทั้งนี้ ANN สามารถเพิ่มความแม่นยำของการพยากรณ์ได้ประมาณ 51.28% เมื่อเปรียบเทียบกับวิธีการเดิมของบริษัท ซึ่งสะท้อนให้เห็นถึงศักยภาพของการเรียนรู้ของเครื่องในการจัดการกับข้อมูลอุปสงค์ที่มีความซับซ้อนและเปลี่ยนแปลงอยู่ตลอดเวลา

อีกประเด็นสำคัญที่งานวิจัยนี้ชี้ให้เห็น คือ ความแม่นยำของการพยากรณ์เพียงอย่างเดียวไม่เพียงพอที่จะสร้างประโยชน์ทางธุรกิจ หากผลการพยากรณ์ไม่ได้ถูกนำไปเชื่อมโยงกับกระบวนการวางแผนในซัพพลายเชน งานวิจัยจึงได้พัฒนาแนวทางที่เชื่อมโยงผลการพยากรณ์เข้ากับการกำหนด

ปริมาณการผลิต ระดับสินค้าสำเร็จรูป และสินค้าระหว่างผลิต เพื่อให้การวางแผนการผลิตและการจัดการสินค้าคงคลังสอดคล้องกับความต้องการของตลาดมากยิ่งขึ้น

ผลลัพธ์ของการประยุกต์ใช้แนวทางดังกล่าวแสดงให้เห็นว่าสามารถลดต้นทุนรวมของระบบได้อย่างมีนัยสำคัญ โดยต้นทุนที่เกิดจากการสูญเสียยอดขายและต้นทุนการถือครองสินค้าคงคลังลดลงประมาณ 1.40 ล้านดอลลาร์สหรัฐ หรือคิดเป็นร้อยละ 27.71 เมื่อเปรียบเทียบกับแนวทางเดิมของบริษัท ผลลัพธ์นี้สะท้อนให้เห็นว่าการประยุกต์ใช้การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องสามารถสร้างคุณค่าทางธุรกิจได้อย่างเป็นรูปธรรม และช่วยเพิ่มประสิทธิภาพของซัพพลายเชนในภาพรวม จากงานวิจัยนี้ สามารถสรุปบทเรียนสำคัญได้ 5 ประการ ดังนี้

1. การพยากรณ์อุปสงค์เป็นรากฐานสำคัญของการวางแผนซัพพลายเชน และส่งผลต่อการตัดสินใจในกิจกรรมต่าง ๆ ตลอดทั้งห่วงโซ่อุปทาน
2. ผลิตภัณฑ์ที่มีวงจรชีวิตสั้นและมีความไม่แน่นอนสูง จำเป็นต้องใช้วิธีการพยากรณ์ที่สามารถเรียนรู้รูปแบบอุปสงค์ที่ซับซ้อนได้อย่างมีประสิทธิภาพ
3. การเรียนรู้ของเครื่อง โดยเฉพาะโครงข่ายประสาทเทียม สามารถช่วยเพิ่มความแม่นยำของการพยากรณ์เมื่อเปรียบเทียบกับวิธีการทางสถิติแบบดั้งเดิม
4. การเชื่อมโยงผลการพยากรณ์เข้ากับการวางแผนการผลิตและการจัดการสินค้าคงคลังเป็นปัจจัยสำคัญที่ช่วยลดต้นทุนและเพิ่มระดับการให้บริการลูกค้า
5. เป้าหมายสูงสุดของการวิเคราะห์ข้อมูลไม่ใช่การสร้างแบบจำลองที่แม่นยำที่สุด แต่คือการสนับสนุนการตัดสินใจที่นำไปสู่ผลลัพธ์ทางธุรกิจที่ดีขึ้น

โดยสรุป งานวิจัยนี้เป็นตัวอย่างที่ชัดเจนของการประยุกต์ใช้การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่องเพื่อแก้ไขปัญหาทางธุรกิจจริงในซัพพลายเชน ตั้งแต่การวิเคราะห์ปัญหา การพัฒนาแบบจำลอง การพยากรณ์ การประเมินประสิทธิภาพของแบบจำลอง ไปจนถึงการนำผลลัพธ์ไปใช้ในการวางแผนการผลิตและการจัดการสินค้าคงคลัง ซึ่งสะท้อนให้เห็นถึงบทบาทของการวิเคราะห์ข้อมูลและ ML ในการสร้างความได้เปรียบทางการแข่งขันและเพิ่มประสิทธิภาพขององค์กรในยุคดิจิทัล

เอกสารอ้างอิง

- Areerakulkan, N., Moryadee, C., Trakoonsanti, L., Khaengkhan, M., & Setthachotsombut, N. (2024). *A new product's demand forecasting using artificial neural network*. In *Proceedings of the 26th International Conference on Enterprise Information Systems (ICEIS 2024)* (Vol. 1, pp. 417–424). SCITEPRESS. <https://doi.org/10.5220/0012734800003690>
- Christopher, M. (2016). *Logistics & supply chain management* (5th ed.). Pearson Education.
- Haykin, S. (2009). *Neural networks and learning machines* (3rd ed.). Pearson Education.
- Ingle, C., Bakliwal, D., Jain, J., Singh, P., Kale, P., & Chhajed, V. (2021). Demand forecasting: Literature review on various methodologies. In *2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT)* (pp. 1–7). IEEE. <https://doi.org/10.1109/ICCCNT51525.2021.9580139>
- Simchi-Levi, D., Kaminsky, P., & Simchi-Levi, E. (2021). *Designing and managing the supply chain: Concepts, strategies, and case studies* (4th ed.). McGraw-Hill Education.
- Vandeput, N. (2021). *Data science for supply chain forecasting* (2nd ed.). Wiley.

บรรณานุกรม

- Alkahtani, M., Al-Ahmari, A., Kaid, H., & Sonboa, M. (2019). Comparison and evaluation of multi-criteria supplier selection approaches: A case study. *Advances in Mechanical Engineering*, 11(2), 1-19. <https://doi.org/10.1177/1687814018822926>.
- Althabatah, A., Yaqot, M., Menezes, B., & Kerbache, L. (2023). Transformative procurement trends: Integrating Industry 4.0 technologies for enhanced procurement processes. *Logistics*, 7(3), Article 63. <https://doi.org/10.3390/logistics7030063>.
- Amindoust, A., Ahmed, S., Saghafein, A., & Bahreininejad, A. (2012). Sustainable supplier selection: A ranking model based on fuzzy inference system. *Applied Soft Computing*, 12(6), 1668-1677. <https://doi.org/10.1016/j.asoc.2012.01.023>.
- Applegate, D. L., Bixby, R. E., Chvátal, V., & Cook, W. J. (2006). *The traveling salesman problem: A computational study*. Princeton University Press.
- Areerakulkan, N., & Pongpech, W. (2021). A Dempster-Shafer big data readiness assessment model. In *Proceedings of the 23rd International Conference on Enterprise Information Systems (ICEIS 2021)* (Vol. 2, pp. 581-585). SCITEPRESS. <https://doi.org/10.5220/0010506205810585>.
- Areerakulkan, N., Lhaocho, L., & Kijwattanaphokin, S. (2026). Vehicle Routing Planning Using a Hybrid Approach of Spectral Clustering and Nearest Neighbour Algorithms. *Journal of the CUFST*, 1(1).
- Areerakulkan, N., Moryadee, C., Trakoonsanti, L., Khaengkhan, M., & Setthachotsombut, N. (2024). A new product's demand forecasting using artificial neural network. In *Proceedings of the 26th International Conference on Enterprise Information Systems (ICEIS 2024)* (Vol. 1, pp. 417-424). SciTePress. <https://doi.org/10.5220/0012734800003690>.
- Bag, S., Choi, T. M., Rahman, M. S., Srivastava, G., & Singh, R. K. (2022). Examining collaborative buyer-supplier relationships and social sustainability in the "new normal" era: the moderating effects of justice and big data analytical intelligence. *Annals of operations research*, 1-46. Advance online publication. <https://doi.org/10.1007/s10479-022-04875-1>.

บรรณานุกรม (ต่อ)

- Bai, R., Chen, X., Chen, Z.-L., Cui, T., Gong, S., He, W., & Zhang, H. (2021). Analytics and machine learning in vehicle routing research: Methods and applications. *International Journal of Production Research*, 59(3), 820-846. <https://doi.org/10.1080/00207543.2021.2013566>.
- Baldi, B., Confente, I., Russo, I., & Gaudenzi, B. (2024). Consumer-centric supply chain management: A literature review, framework, and research agenda. *Journal of Business Logistics*, 45(4), 1-39. <https://doi.org/10.1111/jbl.12399>
- Bandara, K., Bergmeir, C., & Smyl, S. (2020). Forecasting across time series databases using recurrent neural networks on groups of similar series: A clustering approach. *Expert Systems with Applications*, 140, 112896. <https://doi.org/10.1016/j.eswa.2019.112896>
- Baryannis, G., Dani, S., & Antoniou, G. (2019). Predicting supply chain risks using machine learning: The trade-off between performance and interpretability. *Future Generation Computer Systems*, 101, 993-1004. <https://doi.org/10.1016/j.future.2019.07.059>.
- Beamon, B. M. (1999). Measuring supply chain performance. *International Journal of Operations & Production Management*, 19(3), 275-292. <https://doi.org/10.1108/01443579910249714>.
- Bowersox, D. J., Closs, D. J., Cooper, M. B., & Bowersox, J. C. (2019). *Supply chain logistics management* (5th ed.). McGraw-Hill Education.
- Box, G. E. P., Jenkins, G. M., Reinsel, G. C., & Ljung, G. M. (2015). *Time series analysis: Forecasting and control* (5th ed.). Wiley.
- Braekers, K., Ramaekers, K., & Van Nieuwenhuyse, I. (2016). The vehicle routing problem: State of the art classification and review. *Computers & Industrial Engineering*, 99, 300-313. <https://doi.org/10.1016/j.cie.2015.12.007>.
- Bräysy, O., & Gendreau, M. (2005). Vehicle routing problem with time windows, Part I: Route construction and local search algorithms. *Transportation Science*, 39(1), 104-118. <https://doi.org/10.1287/trsc.1030.0056>.

บรรณานุกรม (ต่อ)

- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5-32.
<https://doi.org/10.1023/A:1010933404324>
- Cachon, G. P., & Fisher, M. (2000). Supply chain inventory management and the value of shared information. *Management Science*, 46(8), 1032-1048.
<https://doi.org/10.1287/mnsc.46.8.1032.12029>.
- Cao, L. (2017). Data science: A comprehensive overview. *ACM Computing Surveys*, 50(3), Article 43, 1-42. <https://doi.org/10.1145/3076253>
- Carbonneau, R., Laframboise, K., & Vahidov, R. (2008). Application of machine learning techniques for supply chain demand forecasting. *European Journal of Operational Research*, 184(3), 1140-1154.
<https://doi.org/10.1016/j.ejor.2006.12.004>
- Chaitongrat, S., & Areerakulkan, N. (2021). Application of analytic hierarchy process method for sustainable supplier selection: A case study of ABC Company. *Journal of Logistics and Supply Chain College*, 7(2), 5-17.
- Chanchaichujit, J., Balasubramanian, S., & Charmaine, N. S. M. (2020). A systematic literature review on the benefit-drivers of RFID implementation in supply chains and its impact on organizational competitive advantage. *Cogent Business & Management*, 7(1), 1818408. <https://doi.org/10.1080/23311975.2020.1818408>.
- Chen, D., Sain, S. L., & Guo, K. (2012). Data mining for the online retail industry: A case study of RFM model-based customer segmentation. *Journal of Database Marketing & Customer Strategy Management*, 19(3), 197-208.
<https://doi.org/10.1057/dbm.2012.17>.
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785-794). <https://doi.org/10.1145/2939672.2939785>
- Chen, X., Li, S., & Wang, X. (2020). Evaluating the effects of quality regulations on the pharmaceutical supply chain. *International Journal of Production Economics*, 230, Article 107770. <https://doi.org/10.1016/j.ijpe.2020.107770>.

บรรณานุกรม (ต่อ)

- Choi, T. M., Wallace, S. W., & Wang, Y. (2018). Big data analytics in operations management. *Production and Operations Management*, 27(10), 1868-1883. <https://doi.org/10.1111/poms.12838>
- Chopra, S., & Meindl, P. (2019). *Supply chain management: Strategy, planning, and operation* (7th ed.). Pearson.
- Christopher, M. (2016). *Logistics & supply chain management* (5th ed.). Pearson Education.
- Clarke, G., & Wright, J. W. (1964). Scheduling of vehicles from a central depot to a number of delivery points. *Operations Research*, 12(4), 568-581. <https://doi.org/10.1287/opre.12.4.568>.
- Dantzig, G. B., & Ramser, J. H. (1959). The truck dispatching problem. *Management Science*, 6(1), 80-91. <https://doi.org/10.1287/mnsc.6.1.80>.
- Dantzig, G. B., Fulkerson, D. R., & Johnson, S. M. (1954). Solution of a large-scale traveling-salesman problem. *Journal of the Operations Research Society of America*, 2(4), 393-410. <https://doi.org/10.1287/opre.2.4.393>.
- Dhar, V. (2013). Data science and prediction. *Communications of the ACM*, 56(12), 64-73. <https://doi.org/10.1145/2500499>.
- Dolgui, A., Ivanov, D. and Sokolov, B. (2020) Reconfigurable Supply Chain: The Xnetwork. *International Journal of Production Research*, 58, 4138-4163. <https://doi.org/10.1080/00207543.2020.1774679>.
- Dutta, P., Choi, T.-M., Somani, S., & Butala, R. (2020). Blockchain technology in supply chain operations: Applications, challenges and research opportunities. *Transportation Research Part E: Logistics and Transportation Review*, 142, 102067. <https://doi.org/10.1016/j.tre.2020.102067>.
- Esper, T. L., Ellinger, A. E., Stank, T. P., Flint, D. J., & Moon, M. (2020). Everything old is new again: The age of consumer-centric supply chain management. *Journal of Business Logistics*, 41(4), 286-293. <https://doi.org/10.1111/jbl.12267>.
- Fader, P. S. (2020). *Customer centricity: Focus on the right customers for strategic advantage*. Wharton School Press.

บรรณานุกรม (ต่อ)

- Farahani, R. Z., SteadieSeifi, M., & Asgari, N. (2010). Multiple criteria facility location problems: A survey. *Applied Mathematical Modelling*, 34(7), 1689-1709. <https://doi.org/10.1016/j.apm.2009.10.005>.
- Fedushko, S., & Ustyianovych, T. (2022). E-commerce customers behaviour research using cohort analysis: A case study of COVID-19. *Journal of Open Innovation: Technology, Market, and Complexity*, 8(1), Article 12. <https://doi.org/10.3390/joitmc8010012>.
- Feng, B., & Ye, Q. (2021). Operations management of smart logistics: A literature review and future research. *Frontiers of Engineering Management*, 8, 344-355. <https://doi.org/10.1007/s42524-021-0156-2>.
- Fisher, M. L., & Jaikumar, R. (1981). A generalized assignment heuristic for vehicle routing. *Networks*, 11(2), 109-124. <https://doi.org/10.1002/net.3230110205>.
- Gendreau, M., Laporte, G., & Potvin, J.-Y. (2010). Metaheuristics for the capacitated VRP. In P. Toth & D. Vigo (Eds.), *Vehicle routing: Problems, methods, and applications* (2nd ed., pp. 129-154). SIAM.
- Genuer, R., Poggi, J. M., Tuleau-Malot, C., & Villa-Vialaneix, N. (2015). Random forests for big data. arXiv preprint arXiv:1511.08327. <https://arxiv.org/abs/1511.08327>
- Gillett, B. E., & Miller, L. R. (1974). A heuristic algorithm for the vehicle-dispatch problem. *Operations Research*, 22(2), 340-349. <https://doi.org/10.1287/opre.22.2.340>.
- Gilliland, M. (2010). *The business forecasting deal: Exposing myths, eliminating bad practices, providing practical solutions*. Wiley.
- Govindan, K., & Soleimani, H. (2017). A review of reverse logistics and closed-loop supply chains: A journal of cleaner production focus. *Journal of Cleaner Production*, 142, 371-384. <https://doi.org/10.1016/j.jclepro.2016.03.126>.
- Grant, D. B., Trautrim, A., & Wong, C. Y. (2022). *Sustainable logistics and supply chain management: Principles and practices for sustainable operations and management* (3rd ed.). Kogan Page.

บรรณานุกรม (ต่อ)

- Gunasekaran, A., Patel, C., & Tirtiroglu, E. (2001). Performance measures and metrics in a supply chain environment. *International Journal of Operations & Production Management*, 21(1/2), 71-87. <https://doi.org/10.1108/01443570110358468>.
- Guo, X., Gao, L., Liu, X., & Yin, J. (2017). Improved deep embedded clustering with local structure preservation. In *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence (IJCAI)* (pp. 1753-1759). <https://doi.org/10.24963/ijcai.2017/243>.
- Han, J., Kamber, M., & Pei, J. (2011). *Data mining: Concepts and techniques* (3rd ed.). Elsevier.
- Harrison, A., Skipworth, H., van Hoek, R., & Aitken, J. (2019). *Logistics management and strategy: Competing through the supply chain* (6th ed.). Pearson.
- Haselbeck, F., Kuhn, T., & Dimpler, J. (2022). Machine learning outperforms classical forecasting on horticultural sales predictions. *Smart Agricultural Technology*, 2, 100049. <https://doi.org/10.1016/j.atech.2021.100049>.
- Hastie, T., Tibshirani, R., & Friedman, J. H. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer.
- Haykin, S. (2009). *Neural networks and learning machines* (3rd ed.). Pearson Education.
- Hewamalage, H., Ackermann, K., & Bergmeir, C. (2023). Forecast evaluation for data scientists: Common pitfalls and best practices. *Data Mining and Knowledge Discovery*, 37(2), 788-832. <https://doi.org/10.1007/s10618-022-00894-5>.
- Hillier, F. S., & Lieberman, G. J. (2021). *Introduction to operations research* (11th ed.). McGraw-Hill Education.
- Holden, E., Linnerud, K., & Banister, D. (2014). Sustainable development: Our common future revisited. *Global Environmental Change*, 26, 130-139. <https://doi.org/10.1016/j.gloenvcha.2014.04.006>.
- Hou, Y., Khokhar, M., Zia, S., & Sharma, A. (2022). Assessing the best supplier selection criteria in supply chain management during the COVID-19 pandemic. *Frontiers in Psychology*, 12, 804954. <https://doi.org/10.3389/fpsyg.2021.804954>.

บรรณานุกรม (ต่อ)

- Husnah, M., et al. (2023). Customer segmentation analysis using LRFM-based product and brand loyalty with fuzzy C-means algorithm. In 2023 2nd International Conference for Innovation in Technology (INOCON). <https://doi.org/10.1109/INOCON57975.2023.10100974>.
- Hyndman, R. J., & Athanasopoulos, G. (2018). Forecasting: Principles and practice (2nd ed.). OTexts. <https://otexts.com/fpp2/>.
- Hyndman, R. J., & Athanasopoulos, G. (2021). Forecasting: Principles and practice (3rd ed.). OTexts.
- Hyndman, R. J., Koehler, A. B., Ord, J. K., & Snyder, R. D. (2008). Forecasting with exponential smoothing: The state space approach. Springer.
- Ingle, C., Bakliwal, D., Jain, J., Singh, P., Kale, P., & Chhajed, V. (2021). Demand forecasting: Literature review on various methodologies. In 2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT) (pp. 1–7). IEEE. <https://doi.org/10.1109/ICCCNT51525.2021.9580139>
- Islam, S., & Amin, S. H. (2020). Prediction of probable backorder scenarios in the supply chain using distributed random forest and gradient boosting. Journal of Big Data, 7, 65. <https://doi.org/10.1186/s40537-020-00345-2>
- Ivanov, D. (2021). Supply Chain Viability and the COVID-19 Pandemic: A Conceptual and Formal Generalisation of Four Major Adaptation Strategies. International Journal of Production Research, 59, 3535-3552. <https://doi.org/10.1080/00207543.2021.1890852>.
- Ivanov, D., & Dolgui, A. (2020). A digital supply chain twin for managing the disruption risks and resilience in the era of Industry 4.0. Production Planning & Control, 32(9), 775-788. <https://doi.org/10.1080/09537287.2020.1768450>.
- Ivanov, D., & Dolgui, A. (2020). Viability of intertwined supply networks: Extending the supply chain resilience angles towards survivability. International Journal of Production Research, 58(10), 2904-2915. <https://doi.org/10.1080/00207543.2020.1750727>.

บรรณานุกรม (ต่อ)

- Jain, A. K. (2010). Data clustering: 50 years beyond K-means. *Pattern Recognition Letters*, 31(8), 651-666. <https://doi.org/10.1016/j.patrec.2009.09.011>.
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An introduction to statistical learning: With applications in R* (2nd ed.). Springer.
- Kache, F., & Seuring, S. (2017). Challenges and Opportunities of Digital Information at the Intersection of Big Data Analytics and Supply Chain Management. *International Journal of Operations & Production Management*, 37, 10-36. <https://doi.org/10.1108/ijopm-02-2015-0078>.
- Kamath, R. (2018). Food traceability on blockchain: Walmart's pork and mango pilots with IBM. *The Journal of the British Blockchain Association*, 1(1). [https://doi.org/10.31585/jbba-1-1-\(10\)2018](https://doi.org/10.31585/jbba-1-1-(10)2018).
- Knaflic, C. N. (2015). *Storytelling with data: A data visualization guide for business professionals*. John Wiley & Sons.
- Kochan, C. G., Nowicki, D. R., & Sauser, B. (2018). Impact of cloud-based information sharing on hospital supply chain performance: A system dynamics framework. *International Journal of Production Economics*, 195, 168-185. <https://doi.org/10.1016/j.ijpe.2017.10.008>.
- Kool, W., van Hoof, H., & Welling, M. (2019). Attention, learn to solve routing problems! *arXiv*. <https://doi.org/10.48550/arXiv.1803.08475>.
- Koot, M., Mes, M. R. K., & Iacob, M. E. (2021). A systematic literature review of supply chain decision making supported by the Internet of Things and Big Data Analytics. *Computers & Industrial Engineering*, 154, 107076. <https://doi.org/10.1016/j.cie.2020.107076>.
- Kudyba, S. (2014). *Big data, mining, and analytics: Components of strategic decision making* (1st ed.). Auerbach Publications. <https://doi.org/10.1201/b16666>.
- Kumar, N., et al. (2025). Intelligent customer segmentation: unveiling consumer patterns with machine learning. *Neural Computing and Applications*. <https://doi.org/10.1007/s43995-025-00180-7>.

บรรณานุกรม (ต่อ)

- Kundu, O., James, A. D., & Rigby, J. (2020). Public procurement and innovation: A systematic literature review. *Science and Public Policy*, 47(4), 490-502. <https://doi.org/10.1093/scipol/scaa029>.
- Laporte, G. (2009). Fifty years of vehicle routing. *Transportation Science*, 43(4), 408-416. <https://doi.org/10.1287/trsc.1090.0301>.
- Laporte, G., & Semet, F. (2002). Classical heuristics for the capacitated VRP. In P. Toth & D. Vigo (Eds.), *The vehicle routing problem* (pp. 109-128). Society for Industrial and Applied Mathematics (SIAM). <https://doi.org/10.1137/1.9780898718515.ch5>.
- Lawler, E. L., Lenstra, J. K., Rinnooy Kan, A. H. G., & Shmoys, D. B. (1985). *The traveling salesman problem: A guided tour of combinatorial optimization*. Wiley.
- Lee, H. L., Padmanabhan, V., & Whang, S. (1997). Information distortion in a supply chain: The bullwhip effect. *Management Science*, 43(4), 546-558. <https://doi.org/10.1287/mnsc.43.4.546>.
- Lemon, K. N., & Verhoef, P. C. (2016). Understanding customer experience throughout the customer journey. *Journal of Marketing*, 80(6), 69-96. <https://doi.org/10.1509/jm.15.0420>.
- Lenstra, J. K., & Rinnooy Kan, A. H. G. (1981). Complexity of vehicle routing and scheduling problems. *Networks*, 11(2), 221-227. <https://doi.org/10.1002/net.3230110211>.
- Liu, K. Y. (2022). *Supply Chain Analytics: Concepts, Techniques and Applications*. (1st ed.) Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-92224-5>.
- Lukardi, C., & Hamsal, M. (2020). Warehouse selection using center-of-gravity method in minimizing transportation cost. In *Contemporary Research on Business and Management*. Routledge / Taylor & Francis.
- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2018). Statistical and machine learning forecasting methods: Concerns and ways forward. *PLOS ONE*, 13(3), e0194889. <https://doi.org/10.1371/journal.pone.0194889>

บรรณานุกรม (ต่อ)

- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2018). The M4 competition: Results, findings, conclusion and way forward. *International Journal of Forecasting*, 34(4), 802-808. <https://doi.org/10.1016/j.ijforecast.2018.06.001>.
- Manthou, V., & Vlachopoulou, M. (2001). Bar-code technology for inventory and marketing management systems: A model for its development and implementation. *International Journal of Production Economics*, 71(1-3), 157-164. [https://doi.org/10.1016/S0925-5273\(00\)00115-8](https://doi.org/10.1016/S0925-5273(00)00115-8).
- McFarlane, D., & Sheffi, Y. (2003). The impact of automatic identification on supply chain operations. *The International Journal of Logistics Management*, 14(1), 1-17. <https://doi.org/10.1108/09574090310806503>.
- Mediavilla, M. A., Escudero-Santana, A., & Nieto-Morote, A. (2022). Review and analysis of artificial intelligence methods for demand forecasting in supply chain. *Procedia Computer Science*, 200, 112-121. <https://doi.org/10.1016/j.procir.2022.05.119>
- Memari, A., Dargi, A., Jokar, M. R. A., Ahmad, R., & Rahim, A. R. A. (2019). Sustainable supplier selection: A multi-criteria intuitionistic fuzzy TOPSIS method. *Journal of Manufacturing Systems*, 50, 9-24. <https://doi.org/10.1016/j.jmsy.2018.11.002>.
- Miller, C. E., Tucker, A. W., & Zemlin, R. A. (1960). Integer programming formulation of traveling salesman problems. *Journal of the ACM*, 7(4), 326-329. <https://doi.org/10.1145/321043.321046>.
- Mitchell, T. M. (1997). *Machine learning*. McGraw-Hill.
- Mitra, A., Jain, A., Kishore, A. et al. A Comparative Study of Demand Forecasting Models for a Multi-Channel Retail Company: A Novel Hybrid Machine Learning Approach. *Oper. Res. Forum* 3, 58 (2022). <https://doi.org/10.1007/s43069-022-00166-4>
- Molnar, C. (2022) *Interpretable Machine Learning. A Guide for Making Black Box Models Explainable*, Self-Published. <https://christophm.github.io/interpretable-ml-book/>
- Munzner, T. (2014). *Visualization analysis and design*. CRC Press.

บรรณานุกรม (ต่อ)

- Nazari, M., Oroojlooy, A., Snyder, L. V., & Takáč, M. (2018). Reinforcement learning for solving the vehicle routing problem. arXiv. <https://doi.org/10.48550/arXiv.1802.04240>.
- Neaga, I., Liu, S., Xu, L., Chen, H., & Hao, Y. (2015). Cloud enabled big data business platform for logistics services: A research and development agenda. In Decision support systems V: Big data analytics for decision making (pp. 22-33). Springer. https://doi.org/10.1007/978-3-319-18533-0_3.
- Ngai, E. W. T., Xiu, L., & Chau, D. C. K. (2009). Application of data mining techniques in customer relationship management: A literature review and classification. Expert Systems with Applications, 36(2), 2592-2602. <https://doi.org/10.1016/j.eswa.2008.02.021>.
- Nokeri, T. C. (2021). Forecasting using ARIMA, SARIMA, and the additive model. In T. C. Nokeri (Ed.), Implementing machine learning for finance: A systematic approach to predictive risk and performance analysis for investment portfolios (pp. 21-50). Apress. https://doi.org/10.1007/978-1-4842-7110-0_2.
- Oeser, G., & Romano, P. (2016). An empirical examination of the assumptions of the Square Root Law for inventory centralisation and decentralisation. International Journal of Production Research, 54(8), 2298-2319. <https://doi.org/10.1080/00207543.2015.1071895>.
- Osaba, E., Yang, X.-S., Diaz, F., Lopez-Garcia, P., & Carballedo, R. (2020). An improved discrete bat algorithm for symmetric and asymmetric traveling salesman problems. Engineering Applications of Artificial Intelligence, 48, 59-71. <https://doi.org/10.1016/j.engappai.2015.10.006>.
- Pan, S. J., & Yang, Q. (2010). A survey on transfer learning. IEEE Transactions on Knowledge and Data Engineering, 22(10), 1345-1359. <https://doi.org/10.1109/TKDE.2009.191>.

บรรณานุกรม (ต่อ)

- Pishchulov, G., Trautrim, A., Chesney, T., Gold, S., & Schwab, L. (2019). The voting analytic hierarchy process revisited: A revised method with application to sustainable supplier selection. *International Journal of Production Economics*, 211, 166-179. <https://doi.org/10.1016/j.ijpe.2019.01.025>.
- Provost, F., & Fawcett, T. (2013). *Data science for business: What you need to know about data mining and data-analytic thinking*. O'Reilly Media.
- Purvis, B., Mao, Y., & Robinson, D. (2019). Three pillars of sustainability: In search of conceptual origins. *Sustainability Science*, 14, 681-695. <https://doi.org/10.1007/s11625-018-0627-5>
- Qureshi, M. R. N. M. (2022). Evaluating enterprise resource planning (ERP) implementation for sustainable supply chain management. *Sustainability*, 14(22), 14779. <https://doi.org/10.3390/su142214779>.
- Rahman, A. (2024). Consumer-centric supply chain models for enhancing responsiveness and flexibility in the retail sector. SSRN. <https://doi.org/10.2139/ssrn.5392421>.
- Rajani, R. L., Heggde, G. S., Kumar, R., & Chauhan, P. (2022). Demand management strategies role in sustainability of service industry and impacts performance of company: Using SEM approach. *Journal of Cleaner Production*, 369, 133311. <https://doi.org/10.1016/j.jclepro.2022.133311>
- Razak, G. M., Hendry, L. C., & Stevenson, M. (2023). Supply chain traceability: A review of the benefits and its relationship with supply chain resilience. *Production Planning & Control*, 34(11), 1114-1134. <https://doi.org/10.1080/09537287.2021.1983661>.
- Rodríguez-Espindola, O., Albores, P., & Brewster, C. (2018). Disaster preparedness in humanitarian logistics: A collaborative approach for resource management in floods. *European Journal of Operational Research*, 264(3), 978-993. <https://doi.org/10.1016/j.ejor.2017.01.021>.
- Saaty, T. L. (1980). *The analytic hierarchy process*. McGraw-Hill.

บรรณานุกรม (ต่อ)

- Saaty, T. L. (2016). The analytic hierarchy and analytic network processes for the measurement of intangible criteria and for decision-making. In J. S. Dyer (Ed.), *Multiple criteria decision analysis: State of the art surveys* (2nd ed., pp. 363-419). Springer. https://doi.org/10.1007/978-1-4939-3094-4_10
- Şahin, M. (2021). Location selection by multi-criteria decision-making methods based on objective and subjective weightings. *Knowledge and Information Systems*, 63(8), 1991-2021. <https://doi.org/10.1007/s10115-021-01588-y>.
- Sarkis, J., & Dhavale, D.G. (2015). Supplier selection for sustainable operations: A triple-bottom-line approach using a Bayesian framework. *International Journal of Production Economics*, 166, 177-191. <https://doi.org/10.1016/j.ijpe.2014.11.007>
- Schoenherr, T., & Speier-Pero, C. (2015). Data science, predictive analytics, and big data in supply chain management: Current state and future potential. *Journal of Business Logistics*, 36(1), 120-132. <https://doi.org/10.1111/jbl.12082>.
- Schonlau, M., & Zou, R. Y. (2020). The random forest algorithm for statistical learning. *The Stata Journal*, 20(1), 3-29. <https://doi.org/10.1177/1536867X20909688>
- Scornet, E., Biau, G., & Vert, J. P. (2015). Consistency of random forests. *The Annals of Statistics*, 43(4), 1716-1741. <https://doi.org/10.1214/15-AOS1321>
- Shi, J., & Malik, J. (2000). Normalized cuts and image segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22(8), 888-905. <https://doi.org/10.1109/34.868688>.
- Simchi-Levi, D., Kaminsky, P., & Simchi-Levi, E. (2004). *Managing the supply chain: The definitive guide for the business professional*. McGraw-Hill.
- Simchi-Levi, D., Kaminsky, P., & Simchi-Levi, E. (2021). *Designing and managing the supply chain: Concepts, strategies, and case studies* (4th ed.). McGraw-Hill Education.
- Singh, A., Gutub, A., Nayyar, A., & Khan, M. K. (2023). Redefining food safety traceability system through blockchain: Findings, challenges and open issues. *Multimedia Tools and Applications*, 82(14), 21243-21277. <https://doi.org/10.1007/s11042-022-14006-4>.

บรรณานุกรม (ต่อ)

- Singh, R. K., Chaudhary, N., & Saxena, N. (2018). Selection of warehouse location for a global supply chain: A case study. *IIMB Management Review*, 30(4), 343-354. <https://doi.org/10.1016/j.iimb.2018.08.009>.
- Solomon, M. M. (1987). Algorithms for the vehicle routing and scheduling problems with time window constraints. *Operations Research*, 35(2), 254-265. <https://doi.org/10.1287/opre.35.2.254>.
- Souza, G. C. (2014). Supply chain analytics. *Business Horizons*, 57(5), 595-605. <https://doi.org/10.1016/j.bushor.2014.06.004>.
- Sutton, R. S., & Barto, A. G. (2018). Reinforcement learning: An introduction (2nd ed.). MIT Press.
- Tan, W. C., & Sidhu, M. S. (2022). Review of RFID and IoT integration in supply chain management. *Operations Research Perspectives*, 9, 100229. <https://doi.org/10.1016/j.orp.2022.100229>.
- Tarigan, Z. J. H., Siagian, H., & Jie, F. (2021). Impact of enhanced enterprise resource planning (ERP) on firm performance through green supply chain management. *Sustainability*, 13(8), 4358. <https://doi.org/10.3390/su13084358>.
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, 28(13), 1319-1350. <https://doi.org/10.1002/smj.640>.
- Toth, P., & Vigo, D. (2014). Vehicle Routing: Problems, Methods, and Applications. Society for Industrial and Applied Mathematics. <https://doi.org/10.1137/1.9781611973594>.
- Toth, P., & Vigo, D. (Eds.). (2014). Vehicle routing: Problems, methods, and applications (2nd ed.). SIAM.
- Udenio, M., Vatamidou, E., Fransoo, J. C., & Dellaert, N. (2017). Behavioral causes of the bullwhip effect: An analysis using linear control theory. *IIE Transactions*, 49(10), 980-1000. <https://doi.org/10.1080/24725854.2017.1325026>.

บรรณานุกรม (ต่อ)

- Unhelkar, B., Joshi, S., Sharma, M., Prakash, S., Mani, A. K., & Prasad, M. (2022). Enhancing supply chain performance using RFID technology and decision support systems in the industry 4.0—A systematic literature review. *International Journal of Information Management Data Insights*, 2(2), 100084. <https://doi.org/10.1016/j.jjime.2022.100084>.
- Vandeput, N. (2021). *Data science for supply chain forecasting* (2nd ed.). Wiley.
- Van Steenberghe, R. M., Mes, M. R. K., & Van Harten, A. (2020). Forecasting demand profiles of new products. *Decision Support Systems*, Volume 139, 2020, 113401, <https://doi.org/10.1016/j.dss.2020.113401>
- Von Luxburg, U. (2007). A tutorial on spectral clustering. *Statistics and Computing*, 17(4), 395-416. <https://doi.org/10.1007/s11222-007-9033-z>.
- Waller, M. A., & Fawcett, S. E. (2013). Data science, predictive analytics, and big data: A revolution that will transform supply chain design and management. *Journal of Business Logistics*, 34(2), 77-84. <https://doi.org/10.1111/jbl.12010>.
- Wang, G. (2025). Customer segmentation in the digital marketing using a Q-learning based differential evolution algorithm integrated with K-means clustering. *PLOS ONE* 20(2): e0318519. <https://doi.org/10.1371/journal.pone.0318519>.
- Wang, S., Sun, L., & Yu, Y. (2024). A dynamic customer segmentation approach by combining LRFMS and multivariate time series clustering. *Scientific Reports*, 14, 17491. <https://doi.org/10.1038/s41598-024-68621-2>.
- Wedel, M., & Kannan, P. K. (2016). Marketing analytics for data-rich environments. *Journal of Marketing*, 80(6), 97-121. <https://doi.org/10.1509/jm.15.0413>.
- Xia, B. S., & Gong, P. (2014). Review of business intelligence through data analysis. *Benchmarking*, 21(2), 300-311. <https://doi.org/10.1108/BIJ-08-2012-0050>.
- Zhang, G., Patuwo, B. E., & Hu, M. Y. (1998). Forecasting with artificial neural networks: The state of the art. *International Journal of Forecasting*, 14(1), 35-62. [https://doi.org/10.1016/S0169-2070\(97\)00044-7](https://doi.org/10.1016/S0169-2070(97)00044-7).

บรรณานุกรม (ต่อ)

- Zhuo, Z. (2019). Supply Chain from the Demand Orientation: A Systematic Literature Review and Theoretical Model Construction. *Mathematics and Computer Science*, 4(2), 41-56. <https://doi.org/10.11648/j.mcs.20190402.11>.
- Zhu, X. (2005). Semi-supervised learning literature survey (Technical Report 1530). Department of Computer Sciences, University of Wisconsin-Madison.

ดัชนีท้ายเล่ม

ก

- การจัดกลุ่มเชิงสเปกตรัล (Spectral Clustering), 197-202, 257-273
- การจัดกลุ่มลูกค้า (Customer Segmentation), 43-67, 209-233
 - ด้วย Cohort Analysis, 48-53
 - ด้วย RFM, 53-58, 214-224
 - ด้วย K-means, 58-61, 225-232
 - heatmap เปรียบเทียบผลการจัดกลุ่ม, 231-233
- การจัดการคลังสินค้า, ดู WMS; การจัดการสินค้าคงคลัง
- การจัดการความต้องการสินค้า (Demand Management), 97-100
- การจัดการความสัมพันธ์กับซัพพลายเออร์ (Supplier Relationship Management; SRM), 87-89
- การจัดการซัพพลายเชน (Supply Chain Management; SCM), 11-14
 - หลัก 7Rs, 12-13
 - ความร่วมมือและการแบ่งปันข้อมูล, 13-14
- การจัดการสินค้าคงคลัง (Inventory Management), 15-17, 38-39, 292-293
- การจัดซื้อจัดจ้าง (Procurement), 68-71, 95
 - เปรียบเทียบกับ purchasing, 69-70
 - vertical integration กับ outsourcing, 70-71
- การจัดเส้นทางขนส่ง (Vehicle Routing), 175-208, 253-273
 - ด้วยโปรแกรมเชิงเส้นตรง, 184-190
 - ด้วยวิธีฮิวริสติก, 191-196
- ด้วยการเรียนรู้ของเครื่อง, 197-208, 257-273
- การตัดสินใจ (Decision-Making), 8, 17, 77, 274-276
- การถดถอยเชิงเส้น (Linear Regression; LR), 120-128
 - แบบอย่างง่าย, 122-125
 - แบบพหุคูณ, 126-128
- การติดตามและประเมินผล (Monitoring and Evaluation), 17
- การทำนาย, ดู การพยากรณ์อุปสงค์; การวิเคราะห์เชิงทำนาย
- การแบ่งกลุ่มลูกค้า, ดู การจัดกลุ่มลูกค้า
- การบริหารความเสี่ยง (Risk Management), 90-96
- การบำรุงรักษาเชิงคาดการณ์ (Predictive Maintenance), 38
- การประเมินซัพพลายเออร์ (Supplier Evaluation), 77-89
 - เกณฑ์การประเมิน, 77-80
 - วิธีเชิงลำดับชั้น (AHP), 77-84
- การพยากรณ์ความต้องการสินค้า, ดู การพยากรณ์อุปสงค์
- การพยากรณ์สินค้าใหม่ (New Product Demand Forecasting), 275-301
 - การวิเคราะห์และพัฒนาแบบจำลอง, 276-279
 - การประเมินและเปรียบเทียบผลการพยากรณ์, 280-281
 - การประยุกต์ใช้ผลการพยากรณ์ในซัพพลายเชน, 282-283

- การพยากรณ์อุปสงค์ (Demand Forecasting), 97-147, 234-252, 275-301
 - ความสำคัญต่อซัพพลายเชน, 97-100
 - วิธีอนุกรมเวลา, 100-114
 - วิธี Machine Learning, 120-147, 234-252
 - กรณีศึกษา, 234-252
 - สินค้าใหม่, 275-301
- การวิเคราะห์ข้อมูล (Data Analytics), 1-5, 15-18
 - เชิงพรรณนา, 4, 17
 - เชิงวินิจฉัย, 4, 17
 - เชิงทำนาย, 4, 17
 - เชิงสั่งการ, 4, 17
 - เชิงสำรวจ (EDA), 4-5
- การวิเคราะห์ข้อมูลจัดซื้อจัดจ้าง, 68-96
- การวิเคราะห์ข้อมูลซัพพลายเชน, 1-20
 - องค์ประกอบ, 15-17
 - หลัก SMART, 16-18
- การวิเคราะห์ข้อมูลลูกค้า, 43-67
- การวิเคราะห์ข้อมูลโลจิสติกส์, 148-208, 253-273
- การวิเคราะห์กลุ่มลูกค้า (Cohort Analysis), 48-53, 62-63
- การวิเคราะห์เชิงคาดการณ์ (Predictive Analysis), 4, 17
- การวิเคราะห์เชิงคำแนะนำ (Prescriptive Analysis), 4, 17
- การวิเคราะห์เชิงพรรณนา (Descriptive Analysis), 4, 17
- การวิเคราะห์เชิงวินิจฉัย (Diagnostic Analysis), 4, 17
- การวิเคราะห์อนุกรมเวลา (Time Series Analysis), 100-114
- การออกแบบเครือข่ายโลจิสติกส์ (Logistics Network Design), 156-174
 - การเลือกทำเล, 156-159
 - centralization กับ decentralization, 160-161
 - square root law, 161-163
- การเรียนรู้ของเครื่อง (Machine Learning), 5-7, 120-147, 197-208, 209-273, 275-301
- การเรียนรู้แบบมีผู้สอน (Supervised Learning), 6
- การเรียนรู้แบบไม่มีผู้สอน (Unsupervised Learning), 6
- การเรียนรู้แบบเสริมกำลัง (Reinforcement Learning), 6
- การเรียนรู้แบบโอนย้าย (Transfer Learning), 7
- การเรียนรู้กึ่งควบคุม (Semi-Supervised Learning), 6-7
- กฎรากที่สอง (Square Root Law; SRL), 161-163
- กรณีศึกษา Coca-Cola, 23-24
- กรณีศึกษาการคัดเลือกซัพพลายเออร์อย่างยั่งยืน, 274-288
- กรณีศึกษาการจัดกลุ่มลูกค้าด้วย Machine Learning, 209-233
- กรณีศึกษาการจัดเส้นทางขนส่งแบบผสม, 253-273
- - กรณีศึกษาการพยากรณ์สินค้าใหม่ด้วยโครงข่ายประสาทเทียม, 275-301
 - กรณีศึกษาการพยากรณ์อุปสงค์ด้วย Machine Learning, 234-252

ข

- ข้อมูล (Data), 21-23
- ข้อมูลกึ่งโครงสร้าง (Semi-structured Data), 15, 34
- ข้อมูลขนาดใหญ่ (Big Data), 29-40
 - คุณลักษณะ 5Vs, 33-36
 - การประยุกต์ใช้ในซัพพลายเชน, 36-39
- ข้อมูลเชิงพื้นที่ (Geospatial Data), 255-256, 259-260
- ข้อมูลเชิงประเมิน (Evaluation Data), 276-277
- ข้อมูลดิจิทัล (Digital Data), 23-24
- ข้อมูลที่มีโครงสร้าง (Structured Data), 15, 33
- ข้อมูลที่ไม่มีโครงสร้าง (Unstructured Data), 15, 34

ค

- โครงข่ายประสาทเทียม (Artificial Neural Network; ANN), 135-147, 275-301
 - การพยากรณ์สินค้าใหม่, 275-301
- ความคลาดเคลื่อนของการพยากรณ์, การประเมิน, 110-114, 280-281
- ความสอดคล้อง (Consistency) ใน AHP, 80, 282-285
- ความเสี่ยงของซัพพลายเออร์, 90-96
- ความยั่งยืน, การคัดเลือกซัพพลายเออร์, 274-288
- คลังข้อมูล (Data Warehouse), 3
- คลาวด์คอมพิวติ้ง (Cloud Computing), 9-12, 18-19
 - public cloud, 10-11
 - hybrid cloud, 10-11
 - private cloud, 10-11
 - IaaS, PaaS, SaaS, 11-12

- ค่าดัชนีความสอดคล้อง (Consistency Index; CI), 282-283
- ค่ารากเฉพาะสูงสุด (lambda max), 281-283
- ค่าใช้จ่ายในการได้มาซึ่งลูกค้า (Customer Acquisition Cost; CAC), 54-55
- ค่าเฉลี่ยเคลื่อนที่ (Moving Average; MA), 104, 107

จ

- จุดส่งสินค้า (delivery points), 255-256, 267-272

ช

- ชุดข้อมูลและตัวแปรที่ใช้, 209-211, 234-237, 253-256, 275-277

ด

- ดัชนีท้ายเล่ม, หลักการทำ, ดูเพิ่มเติม Cross-reference; คำย่อ
- แดชบอร์ด (Dashboards), 17
- ดีโป (Depot), 255-256, 259, 261-272
- Distance Matrix, 255-256

ต

- ตารางเปรียบเทียบซัพพลายเชนขับเคลื่อนด้วยข้อมูลกับแบบดั้งเดิม, 26-29
- ตัวชี้วัดความผิดพลาดการพยากรณ์, 110-114, 280-281
- ต้นไม้ตัดสินใจ (Decision Trees), 6

ถ

- ถดถอยพหุคูณ, ดู การถดถอยเชิงเส้น

ท

- ท่าเลที่ตั้ง, การวิเคราะห์, 156-159

ธ

- ธุรกิจที่ขับเคลื่อนด้วยข้อมูล, ดู
 ซัพพลายเชนที่ขับเคลื่อนด้วยข้อมูล

น

- นอร์มัลไลซ์ข้อมูล (Normalization), 79-80, 280-282
- นอร์มัลไลซ์เมตริกซ์เปรียบเทียบ, 79-80, 280-282
- นโยบายตามการแบ่งกลุ่มลูกค้า, 62-64
- Neural Network, ดู โครงข่ายประสาทเทียม
- Nearest Neighbor, 202-208, 257-273

บ

- บล็อกเชน (Blockchain), 34
- บาร์โค้ด (Barcoding System), 30-31
- บูลล์วิปเอฟเฟกต์, ดู ปรากฏการณ์แส้ผ้า

ป

- ปรากฏการณ์แส้ผ้า (Bullwhip Effect), 13-14, 24-25
- ปัญหาการเดินทางของพนักงานขาย (Travelling Salesman Problem; TSP), 185-190
- ปัญหาการจัดเส้นทางขนส่ง (Vehicle Routing Problem; VRP), 175-184
- ปัญญาประดิษฐ์ (Artificial Intelligence; AI), 5, 23-24
- ป่าสุ่ม (Random Forest; RF), 129-134, 239-244
 - ขั้นตอนการสร้างแบบจำลอง, 131-134
- โปรแกรมเชิงเส้นตรง (Linear Programming; LP), 164-174, 184-190

ผ

- ผูกอย่างยั่งยืน, การจัดอันดับ, 285-288

พ

- พยากรณ์อนุกรมเวลา, ดู การวิเคราะห์อนุกรมเวลา
- พิกัดภูมิศาสตร์ (latitude/longitude), 255-256, 259-262
- พื้นที่จัดเก็บสินค้า, 16

ฟ

- ฟังก์ชันต้นทุน, ใน LP, 188-190

ม

- มูลค่าตลอดช่วงชีวิตลูกค้า (Customer Lifetime Value; CLV), 53-55
- มาตรฐาน Saaty, 275, 279
- Mean Absolute Deviation (MAD), 280-281
- Mean Absolute Percentage Error (MAPE), 111-114, 280-281
- Moving Average, ดู ค่าเฉลี่ยเคลื่อนที่
- Machine Learning, ดู การเรียนรู้ของเครื่อง

ร

- Random Forest, ดู ป่าสุ่ม
- Recency-Frequency-Monetary, ดู การวิเคราะห์ RFM
- RFID, 31-32
- RI (Random Index), 283
- RFM, 53-58, 62-64, 214-224
 - การให้คะแนน, 54-58, 218-224
 - การกำหนดกลุ่มลูกค้า, 57-58, 224
 - การเปรียบเทียบกับ K-means, 231-233
- Root Mean Square Error (RMSE), 280-283
- Routing, ดู การจัดเส้นทางขนส่ง

ล

- โลจิสติกส์ (Logistics), 148-174
 - กิจกรรมหลัก, 151-153
 - รูปแบบการขนส่ง, 154-158
- Linear Regression, ดู การถดถอยเชิงเส้น
- Logistics Network Design, ดู การออกแบบเครือข่ายโลจิสติกส์

ว

- วิทยาการข้อมูล (Data Science), 2-8, 18-19
 - องค์ประกอบหลัก, 2-8
- วิธีคลาร์ก-ไรต์ (Clark-Wright Saving Algorithm), 191-196
- วิธีฮิวริสติก (Heuristic), 191-196
- เวกเตอร์น้ำหนัก (weight vector), 282-285

ส

- สินค้าคงคลัง, ดู การจัดการสินค้าคงคลัง
- สหสัมพันธ์ (Correlation Analysis), 4
- สเปกตรัลคลัสเตอริง, ดู การจัดกลุ่มเชิงสเปกตรัล
- สมการเมทริกซ์, ใน AHP, 281-285
- สมาร์ท (SMART), หลักการวิเคราะห์ซัพพลายเชน, 16-18
- สื่อสารผลการวิเคราะห์, 8-10
- เส้นทางขนส่ง, ดู การจัดเส้นทางขนส่ง
- เส้นทางเชิงกลุ่ม (cluster routes), 267-273
- Supply Chain Management, ดู การจัดการซัพพลายเชน
- Supplier Evaluation, ดู การประเมินซัพพลายเออร์
- Supplier Selection, ดู การคัดเลือกซัพพลายเออร์
- Spectral Clustering, ดู การจัดกลุ่มเชิงสเปกตรัล

- สินค้าใหม่, การพยากรณ์อุปสงค์, 275-301

ท

- Holt-Winters (HW), 278-281

อ

- อนุกรมเวลา (Time Series), 100-114
 - รูปแบบข้อมูล, 102-104
 - Moving Average, 104-107
 - ARIMA, 107-109
- อัลกอริทึมเพื่อนบ้านที่ใกล้ที่สุด, ดู Nearest Neighbor
- อัลกอริทึม K-means, 58-61, 225-232
- อินเทอร์เน็ตของสรรพสิ่ง (IoT), 32-33
- อุปสงค์, ดู การพยากรณ์อุปสงค์
- องค์ประกอบของการวิเคราะห์ข้อมูลซัพพลายเชน, 15-17
- เอ็กซ์จีบูสต์ (XGBoost), 248-252
- เอ็กซ์โปเนนเชียลสมูทติง (Exponential Smoothing; ETS), 278-281

คำย่อ / จุดเข้าถึงเสริม

- AHP, ดู กระบวนการวิเคราะห์เชิงลำดับชั้น; การประเมินซัพพลายเออร์
- AI, ดู ปัญญาประดิษฐ์
- ANN, ดู โครงข่ายประสาทเทียม
- ARIMA, ดู อนุกรมเวลา
- CAC, ดู ค่าใช้จ่ายในการได้มาซึ่งลูกค้า
- CI, ดู ค่าดัชนีความสอดคล้อง
- CLV, ดู มูลค่าตลอดช่วงชีวิตลูกค้า
- CR (Consistency Ratio), 282-283
- CPS (Cyber Physical System), 32-33
- CRM, 14, 16, 24, 30

- EDA (Exploratory Data Analysis), 4-5
- EDI (Electronic Data Interchange), 30
- ERP, 14, 16, 24, 31-32
- ETS, ดู เอ็กซ์โปเนนเชียลสมูทติง
- GIS, 32-33
- HW, ดู Holt-Winters
- IoT, ดู อินเทอร์เน็ตของสรรพสิ่ง
- K-means, ดู อัลกอริทึม K-means
- LP, ดู โปรแกรมเชิงเส้นตรง
- MA, ดู ค่าเฉลี่ยเคลื่อนที่
- MAD, ดู Mean Absolute Deviation
- MAPE, ดู Mean Absolute Percentage Error
- RF, ดู ป่าสุ่ม
- RFID, 31-32
- RFM, ดู การวิเคราะห์ RFM
- RI, ดู Random Index
- RMSE, ดู Root Mean Square Error
- RMSPE, 251-252
- SCM, ดู การจัดการซัพพลายเชน
- SMART, ดู สมาร์ท
- SRL, ดู กฎรกที่สอง
- SRM, ดู การจัดการความสัมพันธ์กับซัพพลายเออร์
- TMS, 24, 31, 154
- TSP, ดู ปัญหาการเดินทางของพนักงานขาย
- VRP, ดู ปัญหาการจัดเส้นทางขนส่ง
- WMS, 16, 24, 31, 154
- XGBoost, ดู เอ็กซ์จีบูสต์
- $\lambda_{\max}$, 281-283

การวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่อง สำหรับการจัดการซัพพลายเชน

DATA ANALYTICS AND MACHINE LEARNING FOR
SUPPLY CHAIN MANAGEMENT