

การใช้ระบบแนะนำแบบผสมในการแก้ปัญหาผู้ใช้งานระบบรายใหม่ร่วมกับการใช้
บริบทสำหรับร้านอาหาร

พิมพ์ชนก ปทุมชาติ

วิทยานิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร
วิทยาศาสตรมหาบัณฑิต (การวิเคราะห์ธุรกิจและวิทยาการข้อมูล)
คณะสถิติประยุกต์
สถาบันบัณฑิตพัฒนบริหารศาสตร์

2563

การใช้ระบบแนะนำแบบผสมในการแก้ปัญหาผู้ใช้งานระบบรายใหม่ร่วมกับการใช้
บริบทสำหรับร้านอาหาร
พิมพ์ชนก ปทุมชาติ
คณะสถิติประยุกต์

..... อาจารย์ที่ปรึกษาวิทยานิพนธ์หลัก
(ผู้ช่วยศาสตราจารย์ ดร.วรพล พงษ์เพชร)

คณะกรรมการสอบวิทยานิพนธ์ ได้พิจารณาแล้วเห็นสมควรอนุมัติให้เป็นส่วนหนึ่งของ
การศึกษาตามหลักสูตรวิทยาศาสตรมหาบัณฑิต (การวิเคราะห์ธุรกิจและวิทยาการข้อมูล)

..... ประธานกรรมการ
(ดร.ธนชาติย์ ฤทธิ์บำรุง)

..... กรรมการ
(ผู้ช่วยศาสตราจารย์ ดร.วรพล พงษ์เพชร)

..... กรรมการ
(ผู้ช่วยศาสตราจารย์ ดร.เอกรัฐ รัฐกาญจน์)

..... กรรมการ
(ดร.รัฐศิลป์ รานอกภาณุวัชร)

..... คนบดี
(ผู้ช่วยศาสตราจารย์ ดร.ปราโมทย์ ลีอนาม)

____/____/____

บทคัดย่อ

ชื่อวิทยานิพนธ์	การใช้ระบบแนะนำแบบผสมในการแก้ปัญหาผู้ใช้งานระบบรายใหม่ร่วมกับการใช้บริบทสำหรับร้านอาหาร
ชื่อผู้เขียน	พิมพ์ชนก ปทุมชาติ
ชื่อปริญญา	วิทยาศาสตรมหาบัณฑิต (การวิเคราะห์ธุรกิจและวิทยาการข้อมูล)
ปีการศึกษา	2563

ปัจจุบันระบบแนะนำสินค้าถูกนำมาใช้อย่างแพร่หลายในหลายอุตสาหกรรม โดยจุดประสงค์ของระบบแนะนำคือเพื่อแนะนำสินค้าที่คาดว่าจะชื่นชอบได้ถูกต้อง โดยวิธีการแนะนำสินค้าแบบดั้งเดิมมีหลายวิธี เช่น วิธีการกรองร่วมกัน วิธีการกรองแบบอิงเนื้อหา อย่างไรก็ตามวิธีเหล่านี้ก็ยังมีประสบปัญหาสำคัญบางประการ เช่น ความสามารถในการปรับขนาด ข้อมูลมีปริมาณน้อย หรือผู้ใช้งานระบบรายใหม่ โดยงานวิจัยนี้จะมุ่งเน้นไปที่ปัญหากรณีผู้ใช้งานระบบรายใหม่ ซึ่งได้นำเสนอวิธีแบบผสมในการแก้ปัญหาโดยใช้ข้อมูลประชากรศาสตร์ของลูกค้าและฟังก์ชันอายุร่วมกับวิธีการเรียนรู้ของเครื่องจักรและบริบทอย่างมีเงื่อนไข โดยใช้ข้อมูลจากธุรกิจร้านอาหารหนึ่งในประเทศไทย จากผลการดำเนินงานพบว่าประสิทธิภาพของโมเดลที่นำเสนอมีค่าความถูกต้องในการแนะนำสูงกว่าวิธีการแนะนำแบบสุ่มและแนะนำด้วยเมตริกอันดับที่ร้อยละ 16 และ 10 ตามลำดับ อีกทั้งนำเสนอวิธีการคิดคะแนนความชื่นชอบโดยนัยของลูกค้ารายเก่าในกรณีที่ร้านอาหารไม่มีคะแนนความชื่นชอบจากลูกค้าโดยตรง โดยพบว่าฟังก์ชันการคิดคะแนนความชื่นชอบที่นำเสนอสามารถนำไปประยุกต์ใช้กับข้อมูลของธุรกิจร้านอาหารได้จริงและสามารถแปลงเป็นค่าคะแนนความชื่นชอบโดยนัยที่นำไปใช้ในการทำนายคะแนนความชื่นชอบของลูกค้ารายเก่าในลำดับต่อไปได้

ABSTRACT

Title of Thesis	A Cold-start Hybrid Context Recommendation System for Casual Restaurant
Author	Pimchanok Patumchat
Degree	Master of Science (Business Analytics and Data Science)
Year	2020

Recommendation systems are used in a variety of industries. The goal of these systems is to identify satisfying selections of items for users. Traditional approaches are classified into several methods, such as Collaborative Filtering and Content-based Filtering. However, significant problems such as scalability, data sparsity, and cold-start are still observed. In this paper, we focused on the cold-start problem for new users. We introduced a hybrid-approach by combining demographic and similarity age functions, together with machine learning and conditional contexts from the dataset of a Thailand-based restaurant. The result shows that the proposed model's accuracy is higher than the random recommendation and the top menu recommendation, 16% and 10%, respectively. Moreover, we also discussed the implicit rating function for existing customers with no explicit ratings in this study. We found that it could be applied to the restaurant dataset and converted into an implicit rating which could be used to predict the existing customer ratings.

กิตติกรรมประกาศ

วิทยานิพนธ์ฉบับนี้จะไม่สำเร็จลุล่วงไปได้หากขาดการช่วยเหลือจากอาจารย์ที่ปรึกษา วิทยานิพนธ์ ผู้ช่วยศาสตราจารย์ ดร.วรพล พงษ์เพชร ที่กรุณาให้คำแนะนำปรึกษาตลอดจนปรับปรุงแก้ไขข้อบกพร่องต่าง ๆ ด้วยความเอาใจใส่เป็นอย่างดี ผู้เขียนตระหนักถึงความตั้งใจจริงและความทุ่มเทของอาจารย์ และขอกราบขอบพระคุณเป็นอย่างสูงไว้ ณ ที่นี้

ขอขอบคุณประธานกรรมการการสอบวิทยานิพนธ์ ดร.ธนชาตย์ ฤทธิ์บำรุง รวมไปถึง คณะกรรมการสอบวิทยานิพนธ์ ผู้ช่วยศาสตราจารย์ ดร.เอกรัฐ รัฐกาญจน์ และอาจารย์ผู้ทรงคุณวุฒิจากภายนอก ดร.รัฐศิลป์ รานอกภานุวัชร ที่สละเวลาและให้ข้อเสนอแนะเพื่อปรับปรุงแก้ไขวิทยานิพนธ์ให้อยู่ในทิศทางที่ชัดเจนมากขึ้น

นอกจากนี้ยังขอขอบคุณคณะอาจารย์ทุกท่านที่เคยสั่งสมให้ความรู้มาตลอดการศึกษา ขอขอบคุณเจ้าหน้าที่ประจำคณะสถิติประยุกต์ทุกคนที่คอยให้คำปรึกษา และช่วยเหลือในการศึกษาระดับปริญญาโทมาตลอดการศึกษา

ขอขอบคุณทุนเรียนดีจากสถาบันบัณฑิตพัฒนบริหารศาสตร์ที่มอบทุนสนับสนุนการศึกษาในระดับปริญญาโทตลอด 2 ปี และขอขอบคุณคณะสถิติประยุกต์ที่มอบทุนสนับสนุนในการนำเสนอผลงานวิจัยในงานประชุมวิชาการด้านการวิจัยดำเนินงานแห่งชาติ National Operations Research Network Conference (OR-Net2021) ที่จังหวัดกรุงเทพฯ ประเทศไทยในวันที่ 13 - 14 พฤษภาคม 2564

ขอขอบคุณคุณพ่อ คุณแม่และครอบครัวที่คอยให้การสนับสนุนมาโดยตลอด ขอขอบคุณเพื่อน ๆ ที่เคยช่วยเหลือไม่ว่าทางตรงหรือทางอ้อม และขอบคุณนายภัทรพงษ์ เสวตพงษ์ที่คอยให้คำปรึกษาและช่วยเหลือตลอดมา

ผู้เขียนหวังเป็นอย่างยิ่งว่างานวิจัยฉบับนี้จะมีประโยชน์ไม่มากนักน้อย สำหรับข้อบกพร่องต่าง ๆ ที่อาจจะเกิดขึ้นนั้นผู้เขียนขอน้อมรับผิดเพียงผู้เดียว และยินดีที่จะรับฟังคำแนะนำจากทุกท่านที่ได้เข้ามาศึกษาเพื่อเป็นประโยชน์ในการพัฒนางานวิจัยต่อไป

พิมพ์ชนก ปทุมชาติ

กรกฎาคม 2564

สารบัญ

	หน้า
บทคัดย่อ	ค
ABSTRACT	ง
กิตติกรรมประกาศ.....	จ
สารบัญ.....	ฉ
สารบัญตาราง.....	ณ
สารบัญภาพ	ญ
บทที่ 1 บทนำ	1
1.1 ความเป็นมาและความสำคัญของปัญหา	1
1.2 วัตถุประสงค์ของงานวิจัย.....	3
1.3 นิยามศัพท์เฉพาะ.....	3
1.4 ประโยชน์ที่คาดว่าจะได้รับ.....	4
1.4.1 ประโยชน์จากงานวิจัย	4
1.4.2 ประโยชน์ทางธุรกิจ	4
1.5 ขอบเขตการศึกษา	4
บทที่ 2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง	6
2.1 ทฤษฎีที่เกี่ยวข้อง.....	6
2.1.1 วิธีการกรองแบบอิงเนื้อหา (Content-based Filtering)	6
2.1.2 วิธีการกรองร่วมกัน (Collaborative Filtering).....	7
2.1.3 วิธีแบบผสม (Hybrid Approach).....	8
2.1.4 วิธีวัดประสิทธิภาพของระบบแนะนำ	9
2.2 งานวิจัยที่เกี่ยวข้อง	11

2.2.1 คะแนนความชื่นชอบโดยนัย (Implicit Ratings).....	11
2.2.2 บริบท (Context).....	12
2.2.3 ปัญหาผู้ใช้งานระบบใหม่ (Cold-start).....	13
2.2.4 โครงข่ายประสาทเทียม (Artificial Neural Network).....	14
บทที่ 3 วิธีการดำเนินการวิจัย	15
3.1 การเก็บรวบรวมและจัดเตรียมข้อมูล (Dataset)	15
3.1.1 การจัดเตรียมข้อมูล (Data Preprocessing).....	15
3.1.2 การแสดงข้อมูล (Data Visualization)	17
3.2 ระบบแนะนำสินค้าแบบผสม (A Hybrid Recommendation System).....	19
3.2.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่.....	19
3.2.1.1 แบ่งกลุ่มลูกค้ารายเก่า.....	21
3.2.1.2 แนะนำสินค้าสำหรับลูกค้ารายใหม่.....	21
3.2.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า.....	23
3.2.2.1 ฟังก์ชันการคิดคะแนนความชื่นชอบโดยนัย (Implicit Rating Function) ของ ลูกค้ารายเก่า	24
3.2.2.2 แนะนำสินค้าสำหรับลูกค้ารายเก่า.....	24
บทที่ 4 ผลการดำเนินการวิจัย.....	27
4.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่.....	27
4.1.1 การแบ่งกลุ่มลูกค้ารายเก่า.....	27
4.1.2 โมเดลต้นไม้ (Decision Tree).....	31
4.1.3 โมเดลแนะนำสินค้าสำหรับลูกค้ารายใหม่.....	32
4.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า.....	34
บทที่ 5 สรุป อภิปรายผล และข้อเสนอแนะ	37
5.1 สรุปและอภิปรายผล.....	37

5.1.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่.....	37
5.1.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า.....	38
5.2 ข้อจำกัดและปัญหาที่พบในงานวิจัย.....	39
5.2.1 ข้อจำกัดด้านข้อมูล.....	39
5.2.2 ข้อจำกัดของระบบแนะนำสินค้า.....	39
5.3 ข้อเสนอแนะในอนาคต	39
5.3.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่.....	39
5.3.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า.....	40
บรรณานุกรม.....	41
ประวัติผู้เขียน.....	45

สารบัญตาราง

	หน้า
ตารางที่ 2.1 ตาราง Confusion Matrix	10
ตารางที่ 3.1 ตัวอย่างข้อมูลนำเข้าหลังจากทำ Data Cleaning	16
ตารางที่ 3.2 ตัวอย่างตารางการทำนายกลุ่มให้แก่ลูกค้ารายใหม่ด้วยโมเดลต้นไม้	22
ตารางที่ 4.1 ตาราง Confusion Matrix ของโมเดลต้นไม้	32
ตารางที่ 4.2 ผลลัพธ์ค่าความถูกต้องของโมเดลต้นไม้	32
ตารางที่ 4.3 ตารางเปรียบเทียบประสิทธิภาพการทำนายทั้ง 4 วิธี	33
ตารางที่ 4.4 ตาราง Confusion Matrix ของโมเดลแนะนำรายการอาหาร	34
ตารางที่ 4.5 ประสิทธิภาพการทำนาย Rating ของแต่ละโมเดล	36

สารบัญภาพ

	หน้า
ภาพที่ 3.1 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรเพศ.....	17
ภาพที่ 3.2 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรอายุ.....	17
ภาพที่ 3.3 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรอาชีพ.....	18
ภาพที่ 3.4 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรภูมิลำเนา.....	18
ภาพที่ 3.5 สถาปัตยกรรมโครงสร้างของโมเดล Artificial Neural Network.....	25
ภาพที่ 3.6 กรอบแนวคิดของระบบแนะนำสินค้าแบบผสม.....	26
ภาพที่ 4.1 แสดงกราฟ Elbow Method.....	27
ภาพที่ 4.2 เปรียบเทียบจำนวนลูกค้าทั้ง 7 กลุ่ม.....	28
ภาพที่ 4.3 เปรียบเทียบจำนวนตัวแปรเพศของลูกค้าทั้ง 7 กลุ่ม.....	28
ภาพที่ 4.4 เปรียบเทียบจำนวนตัวแปรอายุของลูกค้าทั้ง 7 กลุ่ม.....	29
ภาพที่ 4.5 เปรียบเทียบจำนวนตัวแปรอาชีพของลูกค้าทั้ง 7 กลุ่ม.....	30
ภาพที่ 4.6 เปรียบเทียบจำนวนตัวแปรภูมิลำเนาของลูกค้าทั้ง 7 กลุ่ม.....	30
ภาพที่ 4.7 กราฟแสดงการเปรียบเทียบประสิทธิภาพของโมเดลทั้ง 2 กรณี.....	34

บทที่ 1

บทนำ

1.1 ความเป็นมาและความสำคัญของปัญหา

ธุรกิจร้านอาหารเป็นหนึ่งในธุรกิจที่มีความสำคัญอย่างยิ่งต่อภาคบริการและเศรษฐกิจของประเทศไทยเนื่องด้วยมีมูลค่าหมุนเวียนที่ไม่ต่ำกว่า 4 แสนล้านบาทต่อปี ประกอบกับมีผู้ประกอบการที่หลากหลายทั้งผู้ประกอบการรายเล็กไปจนถึงผู้ประกอบการรายใหญ่ที่สนใจเข้ามาเริ่มทำธุรกิจในตลาดอย่างต่อเนื่องรวมถึงยังมีความเกี่ยวเนื่องไปยังอุตสาหกรรมอื่น ๆ แม้ว่าในปัจจุบันเศรษฐกิจของประเทศไทยจะมีภาวะผันผวนแต่สำหรับธุรกิจร้านอาหารยังมีแนวโน้มเติบโตอย่างต่อเนื่องราวร้อยละ 4 - 5 ในช่วงปี พ.ศ. 2562 - 2563 โดยในปี พ.ศ. 2561 ที่ผ่านมามีจำนวนร้านอาหารเปิดใหม่ทั้งสิ้น 34,934 ร้าน ซึ่งเพิ่มสูงขึ้นกว่าในปีก่อนคิดเป็นร้อยละ 9.6 โดยปัจจัยสนับสนุนสำคัญมาจากการเปลี่ยนแปลงโครงสร้างประชากรครัวเรือนของคนไทยที่มีขนาดเล็กลงและต้องการความสะดวกสบายมากขึ้น รวมถึงการขยายตัวของเมืองที่เกิดขึ้นพร้อมกับการขยายตัวของศูนย์การค้าใหม่ ๆ ส่งผลให้ผู้บริโภคมีการปรับเปลี่ยนพฤติกรรมหันมาทานอาหารนอกบ้านหรือซื้ออาหารสำเร็จรูปมาทานที่บ้านมากขึ้น ด้วยจำนวนคู่แข่งที่เพิ่มสูงขึ้นส่งผลให้ผู้ประกอบการจำเป็นต้องปรับตัวและหันมาให้ความสำคัญกับการวางแผนการดำเนินธุรกิจมากยิ่งขึ้น อาทิเช่น กลยุทธ์ทางการตลาด หรือกลยุทธ์การบริหารลูกค้าสัมพันธ์ (Customer Relationship Management) ซึ่งการบริการที่น่าพึงพอใจถือเป็นหนึ่งในกลยุทธ์ที่จะทำให้ลูกค้าเกิดความประทับใจและเลือกกลับมาใช้บริการอีกครั้ง

เครื่องมือหนึ่งที่ถูกนำมาช่วยเพิ่มความพึงพอใจและสร้างประสบการณ์ที่ดีในการได้รับบริการของลูกค้าคือ “ระบบแนะนำ” Recommendation System (RS) โดยระบบแนะนำจะทำหน้าที่ช่วยคัดกรองวัตถุ (Item) ที่ตรงกับความต้องการของผู้ใช้จากวัตถุที่มีอยู่เป็นจำนวนมาก วัตถุในความหมายของระบบแนะนำหมายถึงสิ่งที่ต้องการจะแนะนำ เช่น รายการอาหาร, หนังสือ, ภาพยนตร์, บริการ, ข้อมูลข่าวสาร เป็นต้น โดยทั่วไปวิธีการพื้นฐานที่นิยมใช้กันอย่างกว้างขวางในระบบแนะนำแบบเดิมจำแนกได้เป็น 2 วิธีคือ Content-based และ Collaborative Filtering วิธีแนะนำแบบ Content-based จะประเมินค่าความคล้ายคลึงกันระหว่างวัตถุที่จะแนะนำให้กับผู้ใช้กับวัตถุที่ผู้ใช้เคยและชอบในอดีต โดยพิจารณาเนื้อหาหรือข้อมูลที่อธิบายลักษณะของวัตถุนั้น ๆ และ

เลือกที่จะแนะนำวัตถุที่มีลักษณะคล้ายคลึงกับวัตถุที่ผู้ใช้คุ้นเคยและชอบมากในอดีต ส่วนวิธีแนะนำแบบ Collaborative Filtering จะประเมินค่าความคล้ายคลึงกันระหว่างผู้ใช้เป้าหมาย (Target User) กับผู้ใช้คนอื่น ๆ (Neighbors) ในระบบและเลือกที่จะแนะนำวัตถุที่ได้รับคะแนนความพึงพอใจ (Rating) จากผู้ใช้คนอื่นที่มีลักษณะคล้ายกับผู้ใช้เป้าหมายมากที่สุด

โดยระบบแนะนำจะประกอบไปด้วย 3 ส่วนหลักดังนี้

1) ส่วนข้อมูลพื้นฐานที่จำเป็นต้องใช้ในการประมวลผล เช่น Profile ของผู้ใช้งานแต่ละคน เป็นต้น

2) ส่วนการตอบกลับของข้อมูลซึ่งเป็นข้อมูลที่ได้จากการตอบกลับของผู้ใช้ เช่น การให้คะแนน Rating ซึ่งมีอยู่ 2 แบบคือการตอบกลับแบบโดยตรง (Explicit Rating) เช่น ระดับความชื่นชอบตั้งแต่ 1 ถึง 5 เป็นต้น และการตอบกลับแบบโดยนัย (Implicit Rating) เช่น ประวัติการซื้อสินค้า

3) ส่วนของอัลกอริทึมเป็นส่วนที่ใช้ประมวลผลข้อมูลเพื่อแนะนำข้อมูลให้ตรงกับความต้องการของผู้ใช้งานมากที่สุด ซึ่งโดยทั่วไปจะมีอยู่ 3 วิธีหลักคือ Content-Based Recommendation, Collaborative Filtering Recommendation และ Hybrid Approaches

วิธีการแนะนำแบบดั้งเดิมต่าง ๆ ที่กล่าวมานั้นล้วนมีทั้งข้อดีและข้อเสียที่แตกต่างกัน แต่วิธีการแนะนำเหล่านี้ก็ยังประสบปัญหาที่สำคัญบางประการ เช่น ความสามารถในการปรับขนาด (Scalability) ข้อมูลมีปริมาณน้อย (Data Sparsity) หรือปัญหาผู้ใช้งานระบบรายใหม่ (Cold-start Problem) โดยปัญหาผู้ใช้งานระบบรายใหม่สามารถแบ่งออกเป็น 3 ประเภทดังนี้ Parhi, Pal, and Aggarwal (2017)

1) User Cold-start: คือกรณีที่ลูกค้าใหม่เพิ่งเข้าสู่ระบบ ระบบจึงยังไม่มีข้อมูลความชื่นชอบในอดีตของลูกค้า จึงทำให้ไม่สามารถแนะนำสินค้าที่คาดว่าจะลูกค้ารายใหม่จะชื่นชอบได้

2) Item Cold-start: คือกรณีที่สินค้าชิ้นใหม่เพิ่งเข้าสู่ระบบ จึงทำให้สินค้าชิ้นนั้นยังไม่เคยได้รับคะแนนความชื่นชอบจากผู้ใช้งานใดในระบบมาก่อน

3) User-item Cold-start: คือกรณีที่ลูกค้าใหม่และสินค้าใหม่เข้าสู่ระบบ

สำหรับงานวิจัยนี้ผู้วิจัยต้องการศึกษาแนวทางแก้ปัญหาผู้ใช้งานระบบรายใหม่โดยจะมุ่งเน้นศึกษาเฉพาะกรณีการแนะนำสำหรับลูกค้ารายใหม่ ซึ่งเป็นปัญหาเมื่อลูกค้ารายใหม่ที่เพิ่งเข้าสู่ระบบทำให้ระบบไม่สามารถแนะนำสินค้าที่คาดว่าจะลูกค้ารายใหม่จะชื่นชอบได้เนื่องจากระบบยังไม่มีข้อมูลความชื่นชอบของลูกค้าในอดีต โดยแม้จะมีงานวิจัยที่พยายามแก้ปัญหาการแนะนำสินค้าสำหรับลูกค้ารายใหม่หลากหลายวิธี เช่น การนำข้อมูลทางสื่อสังคมออนไลน์หรือข้อมูลด้านบุคลิกภาพของลูกค้ามาช่วยใช้ในการแก้ปัญหา แต่วิธีเหล่านั้นก็ล้วนมีข้อจำกัดและความเหมาะสมที่จะนำไปใช้ของแต่ละ

องค์กรแตกต่างกัน อาทิเช่น ปัญหาความเป็นส่วนตัวของข้อมูล หรือบางวิธีก็อาจไม่เหมาะสมหรือใช้
ได้ผลกับลักษณะข้อมูลที่มีอยู่ของธุรกิจใดธุรกิจหนึ่งเท่านั้น

ไม่เพียงแต่ในกรณีของลูกค้ารายใหม่เท่านั้น หากระบบไม่มีการเก็บข้อมูลคะแนนความชื่นชอบโดยตรง (Explicit Rating) ในอดีตของลูกค้ารายเก่าก็จะก่อให้เกิดปัญหาเช่นกัน เนื่องจากหากไม่มีข้อมูลคะแนนความชื่นชอบของลูกค้า ระบบก็จะไม่มีข้อมูลที่จะนำมาวิเคราะห์ความชื่นชอบของลูกค้าได้ โดยแม้จะมีงานวิจัยในอดีตที่พยายามนำเสนอวิธีการแก้ปัญหาโดยใช้การอนุมานความชื่นชอบสิ่งที่คาดว่าลูกค้าจะสนใจจากสิ่งอื่น ๆ แทนการตอบกลับแบบชัดเจน (Explicit Feedback) ของลูกค้า เช่น การอนุมานจากพฤติกรรมของลูกค้าโดยดูจากประวัติการสั่งซื้อสินค้า พฤติกรรมการค้นหาสินค้า หรือระยะเวลาที่ลูกค้าใช้บนหน้าเว็บไซต์ร้านค้า ซึ่งสิ่งเหล่านี้ถือเป็นการตอบกลับแบบเป็นนัย (Implicit Feedback) แต่วิธีที่เหมาะสมกับธุรกิจหนึ่งก็อาจจะไม่เหมาะที่จะนำมาใช้กับอีกธุรกิจหนึ่งเสมอไป

ดังนั้นจากที่กล่าวมาข้างต้นจึงสรุปได้ว่าผู้วิจัยต้องการนำเสนอระบบแนะนำสินค้าแบบผสมเพื่อใช้แก้ปัญหาผู้ใช้งานระบบรายใหม่ (Cold-start) กรณีลูกค้ารายใหม่ และนำเสนอวิธีการอนุมานความชื่นชอบโดยนัย (Implicit Rating) ของลูกค้ารายเก่าที่เหมาะสมกับธุรกิจร้านอาหารและข้อมูลที่มีอยู่ในปัจจุบัน เพื่อเป็นแนวทางในการสร้างระบบแนะนำสินค้าสำหรับธุรกิจร้านอาหารหรืออุตสาหกรรมอื่น ๆ ต่อไปในอนาคต

1.2 วัตถุประสงค์ของงานวิจัย

- 1) นำเสนอระบบแนะนำสินค้าแบบผสมเพื่อใช้แก้ปัญหาผู้ใช้งานระบบรายใหม่ (Cold-start) กรณีลูกค้ารายใหม่
- 2) นำเสนอวิธีการอนุมานความชื่นชอบโดยนัย (Implicit Rating) ของลูกค้ารายเก่าที่เหมาะสมกับธุรกิจร้านอาหารและข้อมูลที่มีอยู่ในปัจจุบัน

1.3 นิยามศัพท์เฉพาะ

- 1) ระบบแนะนำสินค้า (Recommendation System) หมายถึง ระบบที่ใช้แนะนำสินค้า (Item) ให้กับผู้ใช้งาน (User) แบบรายบุคคลโดยอ้างอิงจากข้อมูลความชื่นชอบของผู้ใช้ในอดีต
- 2) ปัญหาผู้ใช้งานระบบรายใหม่ (User Cold-start Problem) หมายถึง ปัญหาที่ลูกค้ารายใหม่เพิ่งเข้าสู่ระบบ ระบบจึงยังไม่มีข้อมูลความชื่นชอบในอดีตของลูกค้าทำให้ไม่สามารถแนะนำสินค้าที่คาดว่าลูกค้ารายใหม่จะชื่นชอบได้

- 3) ผู้ใช้ระบบ (User) หมายถึง ผู้ใช้งานระบบแนะนำสินค้า ซึ่งระบบต้องการแนะนำสินค้าที่ตรงกับความต้องการของผู้ใช้มากที่สุด
- 4) สินค้า (Item) หมายถึง วัตถุหรือผลิตภัณฑ์ที่ระบบแนะนำสินค้าต้องการที่จะแนะนำให้แกผู้ใช้ระบบแบบรายบุคคล
- 5) คะแนนความชื่นชอบ (Rating) หมายถึง ระดับความพึงพอใจที่ผู้ใช้แต่ละคนมีต่อสินค้าแต่ละชิ้น
- 6) การตอบกลับแบบโดยตรง (Explicit Rating) หมายถึง การตอบกลับคะแนนความชื่นชอบที่ผู้ใช้ให้แกระบบโดยตรง เช่น การกดถูกใจ (Like) หรือการให้คะแนนความพึงพอใจ (Rating)
- 7) การตอบกลับแบบโดยนัย (Implicit Rating) หมายถึง การอนุมานความชื่นชอบสิ่งที่คาดว่าผู้ใช้จะสนใจจากสิ่งอื่น ๆ แทนการตอบกลับแบบชัดเจน เช่น จากพฤติกรรมการใช้งาน

1.4 ประโยชน์ที่คาดว่าจะได้รับ

1.4.1 ประโยชน์จากงานวิจัย

- 1) ระบบแนะนำสินค้าแบบผสมซึ่งใช้ในการแนะนำสินค้าสำหรับลูกค้ารายใหม่
- 2) วิธีการคิดคะแนนความชื่นชอบโดยนัย (Implicit Rating) ของลูกค้ารายเก่าที่สามารถนำไปประยุกต์ใช้ได้จริง

1.4.2 ประโยชน์ทางธุรกิจ

- 1) เพิ่มยอดขาย เมื่อระบบสามารถรู้ถึงสิ่งที่ลูกค้าต้องการอย่างแท้จริง ก็จะสามารถแนะนำในสิ่งที่เหมาะสมและตรงใจลูกค้าได้มากขึ้น
- 2) เพิ่มความหลากหลายให้แก่ลูกค้า ระบบจะช่วยแนะนำรายการอาหารที่ลูกค้าอาจไม่เคยเลือกซื้อมาก่อนแต่มีแนวโน้มที่จะสนใจในสินค้านั้น
- 3) เพิ่มความพึงพอใจ ช่วยเพิ่มประสบการณ์ที่ดีในการได้รับการแนะนำรายการอาหารที่ลูกค้ามีความสนใจ
- 4) เพิ่มการกระจายคลังสินค้าให้มีความทั่วถึงมากขึ้น

1.5 ขอบเขตการศึกษา

ขอบเขตในการศึกษาครั้งนี้ผู้วิจัยใช้ชุดข้อมูลจริงจากธุรกิจร้านอาหารหนึ่งในประเทศไทย โดยผู้วิจัยใช้ข้อมูลธุรกรรมการสั่งซื้ออาหารตั้งแต่เดือนมกราคมถึงเดือนธันวาคมในปี พ.ศ. 2562

โดยนำมาวิเคราะห์เฉพาะกรณีการสั่งซื้อประเภทอาหารจานหลักของลูกค้าที่เป็นสมาชิกเท่านั้น (ไม่รวมการสั่งซื้ออาหารประเภทเครื่องเคียง ของทานเล่น เครื่องดื่ม และอาหารโปรโมชันตามเทศกาลต่าง ๆ) จากทุกสาขาในประเทศไทยจำนวน 179 สาขา และเนื่องจากเป็นชุดข้อมูลจริงจึงทำให้เกิดข้อจำกัดในบางประการ เช่น การเข้าถึงข้อมูลของลูกค้าและแหล่งข้อมูลที่เกี่ยวข้องมีอยู่อย่างจำกัด รวมถึงลักษณะของข้อมูลที่มีจำนวนว่างค่อนข้างมาก (Sparse Data) ซึ่งเป็นส่วนสำคัญที่อาจส่งผลต่อการวิเคราะห์และค่าความถูกต้องของระบบแนะนำสินค้า



บทที่ 2

ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

เนื้อหาในบทที่ 2 ผู้วิจัยนำเสนอรายละเอียดทฤษฎีและงานวิจัยที่เกี่ยวข้อง พร้อมทั้งสรุปข้อจำกัดของงานวิจัยต่าง ๆ ที่ผ่านมา เพื่อนำมาต่อยอดเป็นแนวทางในการพัฒนาในงานวิทยานิพนธ์ต่อไป โดยรายละเอียดในการนำเสนอจะแบ่งออกเป็น 2 ส่วนหลักดังนี้

2.1 ทฤษฎีที่เกี่ยวข้อง

2.1.1 วิธีการกรองแบบอิงเนื้อหา (Content-based Filtering)

2.1.2 วิธีการกรองร่วมกัน (Collaborative Filtering)

2.1.3 วิธีแบบผสม (Hybrid)

2.1.4 วิธีวัดประสิทธิภาพของระบบแนะนำ

2.2 งานวิจัยที่เกี่ยวข้อง

2.2.1 คะแนนความชื่นชอบโดยนัย (Implicit Ratings)

2.2.2 บริบท (Context)

2.2.3 ปัญหาผู้ใช้งานระบบใหม่ (Cold-start)

2.2.4 โครงข่ายประสาทเทียม (Artificial Neural Network)

2.1 ทฤษฎีที่เกี่ยวข้อง

2.1.1 วิธีการกรองแบบอิงเนื้อหา (Content-based Filtering)

วิธีการกรองแบบอิงเนื้อหา (Content-based Filtering) เป็นวิธีที่แนะนำสิ่งของ (Item) โดยพิจารณาจากคุณสมบัติของสิ่งของที่ผู้ใช้เคยชื่นชอบในอดีต โดยระบบจะแนะนำสิ่งของที่มีคุณสมบัติตรงกับสิ่งที่ผู้ใช้ชื่นชอบในอดีต โดยวิธีการกรองแบบอิงเนื้อหาจะมีข้อแตกต่างจากวิธีการกรองร่วมกันคือ แทนที่ระบบจะแนะนำวัตถุโดยใช้ข้อมูลของผู้ใช้หรือวัตถุที่มีความใกล้เคียงกัน ระบบจะทำการเลือกวัตถุโดยอิงจากความพึงพอใจของผู้ใช้ว่าชอบวัตถุในลักษณะใด คุณสมบัติอย่างไร ซึ่งยังแสดงให้เห็นอีกด้วยว่าผู้ใช้แต่ละคนมีเกณฑ์การเลือก สนใจหรือยึดคุณสมบัติใดของวัตถุเป็นหลัก โดยทั่วไปแล้วระบบแนะนำประเภทนี้จะมีขั้นตอนการทำงานดังนี้

- 1) สร้างประวัติผู้ใช้ระบบโดยอิงจากลักษณะหรือคุณสมบัติของวัตถุซึ่งคำนวณจากประวัติการใช้งานในอดีตของผู้ใช้
- 2) คำนวณความสัมพันธ์ระหว่างประวัติผู้ใช้และวัตถุ
- 3) นำผลลัพธ์ที่ได้จากข้อ 2 มาปรับค่าน้ำหนักแต่ละคุณสมบัติของวัตถุ
- 4) เลือกวัตถุที่มีความสัมพันธ์กับผู้ใช้เป้าหมายมากที่สุด โดยดูจากความสำคัญที่ผู้ใช้เป้าหมายให้ไว้ในแต่ละคุณสมบัติและค่าน้ำหนักของแต่ละคุณสมบัติของวัตถุที่มีอยู่ในระบบ

Content-based เป็นวิธีการหนึ่งของการทำ Information Filtering โดยวิธีการทำงานของเทคนิคนี้คือจะให้ความสนใจเนื้อหาของข้อมูลเป็นสำคัญ เช่น คุณลักษณะ (Feature) เพื่อค้นหาข้อมูลที่ใช้คนนั้นสนใจ ซึ่งวิธีการของเทคนิคดังกล่าวจะไม่ประสบกับปัญหาการให้ Rating ต่อชิ้นข้อมูลที่ไม่มีค่า และปัญหาชิ้นข้อมูลที่ยังไม่ได้ให้ Rating ดังนั้นวิธีการนี้จะต้องมีการคำนวณหาค่าความคล้ายคลึงกัน (Similarity Measure) ระหว่างข้อมูลหรือสินค้าที่ระบบมีอยู่กับความต้องการในตัวข้อมูลหรือสินค้าที่ลูกค้าต้องการ ซึ่ง Algorithm ที่นิยมใช้มากที่สุดคือการทำนายด้วยวิธีการหาสมาชิกที่ใกล้เคียงที่สุด K-NN (K Nearest Neighbor) ซึ่งจะต้องมีการคำนวณหาระยะห่างระหว่างข้อมูล (Standard Euclidean Distance) ข้อดีคือให้ผลลัพธ์ที่มีความถูกต้องสูง เรียนรู้เร็ว และรองรับข้อมูลจำนวนมาก

วิธีการกรองแบบอิงเนื้อหาสามารถแก้ไขปัญหาที่เกิดขึ้นกับวิธีการกรองร่วมกันได้ เนื่องจากในการคำนวณแต่ละครั้งไม่มีการใช้ข้อมูลจากผู้ใช้ที่เคยเลือกวัตถุที่มีความใกล้เคียงกันกับผู้ใช้เป้าหมาย ระบบจึงสามารถให้คำแนะนำได้แม้ว่าจะมีข้อมูลอยู่ในระบบเป็นจำนวนน้อย ในทางกลับกันระบบแบบนี้ก็ยังคงมีปัญหาหลายอย่างเช่น ระบบจะไม่สามารถแนะนำวัตถุอื่น ๆ ที่อยู่นอกเหนือความสนใจของผู้ใช้ระบบได้ เพราะระบบใช้ข้อมูลที่แสดงให้เห็นว่าผู้ใช้ใช้เกณฑ์อะไรในการเลือกวัตถุทำให้ระบบแนะนำเพียงสินค้านั้นแบบเดิม ๆ ซึ่งเป็นการเจาะจงมากเกินไป รวมถึงปัญหาในกรณีที่หากเป็นผู้ใช้ใหม่ในระบบ การแนะนำวิธีนี้จะไม่สามารถแนะนำได้ซึ่งเรียกว่าปัญหาผู้ใช้งานรายใหม่ในระบบ (Cold-start Problem)

2.1.2 วิธีการกรองร่วมกัน (Collaborative Filtering)

การแนะนำสิ่งของด้วยวิธีการกรองร่วมกัน (Collaborative Filtering) เป็นวิธีการแนะนำโดยพิจารณาจากข้อมูลของผู้ใช้ในอดีตเช่นเดียวกับ Content-based Filtering แต่จะต่างกันที่วิธี Collaborative Filtering จะนิยมใช้กับระบบผู้แนะนำที่มีการให้คะแนนความชื่นชอบ (Rating) ต่อสิ่งของด้วย โดยระบบจะสร้าง User Profile เก็บข้อมูลของผู้ใช้แต่ละคนเพื่อใช้ในการคำนวณความคล้ายคลึง ซึ่งความคล้ายคลึงที่คำนวณได้มี 2 ประเภท คือ ความคล้ายคลึงของผู้ใช้ (User-base Similarity) และความคล้ายคลึงของสิ่งของ (Item-base Similarity) เมื่อหาความคล้ายคลึงเรียบร้อยแล้ว

ระบบแนะนำจะคำนวณคะแนนความชื่นชอบที่คาดว่าผู้ใช้เป้าหมาย (Target User) จะให้คะแนนต่อทุก ๆ สิ่งของเพื่อให้ระบบแนะนำแนะนำสิ่งของที่มีค่าทำนายคะแนนความชื่นชอบ (Rating Prediction) สูงที่สุดให้แก่ผู้ใช้เป้าหมาย โดยปกติระบบแนะนำประเภทนี้จะมีขั้นตอนในการประมวลผล 3 ขั้นตอนหลักดังนี้

- 1) สร้างประวัติผู้ใช้หรือวัตถุโดยอิงตามข้อมูลที่จะใช้เป็นพื้นฐานของระบบ
- 2) คัดเลือกผู้ใช้ที่เคยเลือกวัตถุนั้นหรือวัตถุที่เคยถูกเลือกจากผู้ใช้คนดังกล่าวที่มีความใกล้เคียงหรือคล้ายคลึงกันขึ้นมาตามจำนวนที่กำหนดไว้ โดยการเปรียบเทียบประวัติผู้ใช้หรือวัตถุตามแต่ข้อมูลที่ใช้เป็นพื้นฐานของระบบ ซึ่งอาจใช้ค่าความเหมือนโคไซน์ (Cosine Similarity) หรือระยะทางแบบยุคลิด (Euclidean Distance)
- 3) คำนวณว่าผู้ใช้และวัตถุนั้นมีความเหมาะสมกันมากน้อยเพียงใด โดยอาศัยข้อมูลที่หาได้จากขั้นตอนก่อนหน้า

ข้อดีของวิธีนี้คือระบบสามารถแนะนำสิ่งของที่มีความคล้ายคลึงกับสิ่งของในอดีตที่ผู้ใช้ชอบได้โดยไม่ต้องจำเป็นต้องเหมือนกันทั้งหมด และสามารถแนะนำสิ่งของใหม่ ๆ จากผู้ใช้คนอื่น ๆ ในระบบที่มีความคล้ายคลึงกับผู้ใช้เป้าหมาย (User Similarity) ทำให้เพิ่มโอกาสแนะนำสิ่งของประเภทอื่น ๆ ได้มากขึ้น แต่ปัญหาที่พบคือหากเป็นผู้ใช้รายใหม่ในระบบ การแนะนำวิธีนี้จะไม่สามารถแนะนำให้ได้ เช่นเดียวกับกับวิธี Content-based Filtering ซึ่งปัญหานี้เรียกว่า Cold-start Problem

2.1.3 วิธีแบบผสม (Hybrid Approach)

วิธีการแนะนำแบบพื้นฐานต่าง ๆ ที่กล่าวมานั้นมีทั้งข้อดีและข้อเสียที่แตกต่างกันไป โดยนักวิจัยบางท่านได้นำเทคนิควิธีการแนะนำแบบพื้นฐานต่าง ๆ มาผสมผสานกันเพื่อปรับปรุงคุณภาพของการแนะนำให้มีความถูกต้องมากยิ่งขึ้น โดยพยายามรักษาข้อดีและขจัดข้อเสียบางอย่างของวิธีการที่นำมาผสมผสานกันนี้ให้ได้มากที่สุด วิธีการแนะนำในลักษณะนี้เรียกว่าวิธีแบบผสม (Hybrid Approach)

ระบบแนะนำมากมายใช้วิธีการแนะนำข้อมูลแบบผสม (Hybrid) โดยการรวมวิธี Content-Based และ Collaborative Filtering เข้าด้วยกัน ซึ่งช่วยหลีกเลี่ยงข้อจำกัดของวิธี Content-Based และ Collaborative Filtering ได้แต่อาจทำให้มีความซับซ้อนและใช้ทรัพยากรในการแนะนำสูง ซึ่งมีหลายงานวิจัยที่แสดงให้เห็นว่ามีความแม่นยำมากกว่าวิธี Content-Based และ Collaborative Filtering วิธีการที่จะนำวิธี Content-Based และ Collaborative Filtering มาใช้ร่วมกันเป็นระบบแนะนำแบบผสมหรือ Hybrid สามารถแบ่งได้ดังนี้

- 1) Combining Separate Recommenders เป็นทางเลือกหนึ่งในการสร้างระบบแนะนำแบบ Hybrid ซึ่งก็คือการใช้ Content-Based และ Collaborative Filtering แยกกันคนละ

ส่วน จากนั้นจะมีผลลัพธ์ออกมาสองแบบ โดยแบบแรกสามารถรวมผลลัพธ์ที่ได้จากระบบแนะนำแต่ละแบบแล้วให้การแนะนำเดียวกัน โดยจะใช้วิธีการรวมแบบเชิงเส้น (Linear) หรือการลงคะแนน (Vote) หรืออีกทางหนึ่งคือใช้ผลลัพธ์จากระบบแนะนำใดระบบแนะนำหนึ่งเท่านั้น ซึ่งจะเลือกจากผลลัพธ์ที่ดีที่สุดโดยอ้างอิงจากเมตริกคุณภาพในการแนะนำ

2) Adding Content-Based Characteristics to Collaborative Models เป็นการนำคุณสมบัติเฉพาะบางอย่างของวิธี Content-Based มาใส่ไว้ในวิธี Collaborative Filtering ซึ่งมีหลายระบบที่ทำระบบแนะนำแบบ Hybrid ในรูปแบบนี้ เช่น Fab หรือ Collaborative Via Content ซึ่งใช้ Collaborative แบบดั้งเดิม ข้อมูลความชอบของผู้ใช้เป็นแบบ Content-Based ซึ่งเป็นการแก้ปัญหาของข้อมูลของผู้ใช้ที่มีการให้คะแนนวัตถุน้อยเกินไปซึ่งทำให้เกิดข้อจำกัดเรื่อง Data Sparsity

3) Adding Collaborative Characteristics to Content-Based Models เป็นวิธีที่ทำตรงกันข้ามกับวิธีก่อนหน้านี้ โดยจะนำคุณสมบัติเฉพาะบางอย่างของวิธี Collaborative Filtering มาใส่ไว้ในวิธี Content-Based วิธีที่นิยมใช้กันคือการใช้เทคนิคการลดมิติของกลุ่มข้อมูลผู้ใช้ที่อ้างอิงตามเนื้อหา

4) Developing a Single Unifying Recommendation Model เป็นการสร้างโมเดลทางคณิตศาสตร์จากลักษณะเฉพาะของวิธี Content-Based และ Collaborative Filtering เช่น Bayesian Mixed-Effects Regression Models ซึ่งใช้วิธีการของ Markov Chain Monte Carlo ในการประมาณค่าพารามิเตอร์แล้วคาดคะเนหาค่า Rating ที่ไม่ทราบค่า

2.1.4 วิธีวัดประสิทธิภาพของระบบแนะนำ

วิธีวัดประสิทธิภาพที่นิยมใช้ในงานวิจัยต่าง ๆ เพื่อทดสอบผลลัพธ์ที่ได้จากวิธีที่ใช้ในการแนะนำสิ่งของมี 3 รูปแบบคือ การวัดค่าความถูกต้องของระบบแบบภาพรวม (Accuracy), การวัดค่าความแม่นยำของการแนะนำสิ่งของ (Precision) และการวัดค่าความระลึกได้ของการแนะนำสิ่งของ (Recall) โดยนำผลลัพธ์ที่ได้จากระบบแนะนำมาสร้างเป็น Confusion Matrix ประกอบด้วย 4 ช่องคือ True Positive (TP), True Negative (TN), False Positive (FP) และ False Negative (FN)

True Positive (TP) คือระบบแนะนำแนะนำสินค้าขึ้นเดียวกับสินค้าที่ผู้ใช้เป้าหมายชื่นชอบ

True Negative (TN) คือระบบแนะนำไม่ได้แนะนำสินค้าขึ้นนี้ และในความเป็นจริงผู้ใช้เป้าหมายก็ไม่ได้ชอบสินค้าขึ้นดังกล่าวด้วยเช่นกัน

False Positive (FP) คือระบบแนะนำแนะนำสินค้าขึ้นนี้ให้แก่ผู้ใช้เป้าหมาย แต่ความจริงแล้วผู้ใช้เป้าหมายไม่ได้ชอบสินค้าขึ้นดังกล่าว

False Negative (FN) คือระบบแนะนำไม่ได้แนะนำสินค้าชิ้นนี้ให้แก่ผู้ใช้เป้าหมาย แต่ความจริงแล้วผู้ใช้เป้าหมายชอบสินค้าชิ้นดังกล่าว

ตารางที่ 2.1 ตาราง Confusion Matrix

		ค่าที่โมเดลทำนาย	
		Positive	Negative
ค่าจริง	Positive	TP	FN
	Negative	FP	TN

2.1.4.1 ค่าความถูกต้อง (Accuracy)

การวัดค่าความถูกต้องคือการตรวจสอบว่าผลลัพธ์ที่ได้จากระบบแนะนำตรงกับ ความชื่นชอบของผู้ใช้เป้าหมายหรือไม่

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FN + FP}$$

2.1.4.2 ค่าความแม่นยำ (Precision)

การวัดค่าความแม่นยำคือการตรวจสอบว่าผลลัพธ์ที่ได้จากระบบทั้งหมดตรงกับ ความชื่นชอบของผู้ใช้จริงเท่าใด คำนวณจากสัดส่วนของสินค้าที่ระบบแนะนำและตรงกับที่ผู้ใช้ เป้าหมายชื่นชอบต่อสินค้าที่ระบบแนะนำทั้งหมด

$$\text{Precision} = \frac{TP}{TP + FP}$$

2.1.4.3 ค่าความระลึกได้ (Recall)

ค่าความระลึกได้คือการตรวจสอบว่าข้อมูลที่แนะนำจากระบบตรงกับ ความชื่นชอบของผู้ใช้เป้าหมายเท่าใด คำนวณจากสัดส่วนของสินค้าที่ระบบแนะนำและตรงกับที่ผู้ใช้เป้าหมายชื่นชอบต่อสินค้าที่ผู้ใช้เป้าหมายชื่นชอบจริงทั้งหมด

$$\text{Recall} = \frac{TP}{TP + FN}$$

2.2 งานวิจัยที่เกี่ยวข้อง

2.2.1 คะแนนความชื่นชอบโดยนัย (Implicit Ratings)

แบ่งเป็น 2 หัวข้อในการทบทวนวรรณกรรม ได้แก่การอนุมานความชื่นชอบของผู้ใช้ระบบ โดยดูจากข้อมูลการซื้อสินค้า และการอนุมานความชื่นชอบของผู้ใช้ระบบโดยดูจากพฤติกรรมของผู้ใช้ระบบ

2.2.1.1 การอนุมานความชื่นชอบของผู้ใช้ระบบโดยดูจากข้อมูลการซื้อสินค้า

Lee, Park, and Park (2008) ใช้ข้อมูลการซื้อสินค้าของผู้ใช้ร่วมกับข้อมูลเวลา 2 ชนิดคือ 1.ระยะเวลาที่ผู้ใช้ซื้อสินค้าชิ้นนั้น 2.ระยะเวลาที่สินค้าชิ้นนั้นถูกปล่อยออกมา นำมาใช้ในการสร้าง Rating Matrix โดยแต่ละตัวแปรเวลาจะแบ่งออกเป็น 3 ช่วงเวลาได้แก่ ผู้ใช้ซื้อสินค้ามาแล้วซื้อซ้ำกี่ครั้ง และซื้อเมื่อเร็ว ๆ นี้ เช่นเดียวกันกับตัวแปรระยะเวลาที่สินค้าถูกปล่อยออกมาแบ่งเป็น ปล่อยสินค้านานแล้ว ปล่อยสินค้ามาถี่ๆ และเพิ่งปล่อยสินค้าเมื่อเร็ว ๆ นี้ โดย Rating Function จะทำการ Weight คำนวณน้ำหนักจากระยะเวลาที่ผู้ใช้แต่ละคนซื้อสินค้าแต่ละชิ้นร่วมกับระยะเวลาที่สินค้าชิ้นนั้นถูกปล่อยออกมา โดยมีข้อกำหนดว่าหากผู้ใช้เพิ่งซื้อสินค้าเมื่อเร็ว ๆ นี้จะได้รับค่าน้ำหนักคะแนนที่สูงกว่าประเภทอื่น และสินค้าที่เพิ่งปล่อยออกมาก็จะมีคะแนนสูงกว่าสินค้าประเภทอื่น ๆ เช่นเดียวกัน

2.2.1.2 การอนุมานความชื่นชอบของผู้ใช้ระบบโดยดูจากพฤติกรรมของผู้ใช้ระบบ

งานวิจัยของ Covington, Adams, and Sargin (2016) ทำการอนุมานความชื่นชอบของผู้ใช้ระบบโดยดูจากพฤติกรรมการดู Video บน Youtube ของผู้ใช้แทนการตอบกลับแบบชัดเจนหรือการกดถูกใจ (Like) ซึ่งมีข้อมูลว่างค่อนข้างมาก (Sparse Data) โดยระบบจะทำการอนุมานว่าผู้ชมมีความชื่นชอบคลิปวิดีโอหากผู้ใช้ดูคลิปวิดีโอจนจบ เช่นเดียวกันกับงานวิจัยของ Kim and Chan (2005) ที่จะอนุมานว่าผู้ใช้ชื่นชอบเว็บเพจ หากผู้ใช้เลื่อนหน้าเว็บเพจจนถึงด้านล่างสุด ส่วนงานวิจัยของ Haoxian (2017) ซึ่งสร้างระบบแนะนำร้านอาหารออนไลน์ได้ใช้ข้อมูล Log File ของผู้ใช้ระบบแทนคะแนนความชื่นชอบแบบโดยตรง (Explicit Rating) โดยดูจากพฤติกรรมของผู้ใช้บนหน้าเว็บไซต์ร้านค้า โดยกำหนดให้แต่ละพฤติกรรมมีคะแนนในช่วง -1 ถึง 1 ซึ่งพฤติกรรมที่มีค่าเข้าใกล้ 1 จะแสดงถึงผู้ชมมีแนวโน้มที่จะชื่นชอบร้านอาหารที่กำลังค้นหาในปัจจุบัน เช่นการคลิกปุ่มออกจากการค้นหา (Exit Point) จะให้คะแนนความชื่นชอบเท่ากับ 0.8 เนื่องจากการอนุมานว่าผู้ชมมีแนวโน้มที่จะชื่นชอบร้านอาหารปัจจุบันแล้วจึงหยุดการค้นหาต่อ แต่ก็ไม่สามารถสรุปได้อย่างแน่นอน เนื่องจากการคลิกปุ่มออกจากการค้นหาอาจทำตามขั้นตอนทั่วไปเท่านั้น หรือพฤติกรรมอื่นๆ เช่นการกดค้นหา (Browse) ซึ่งจะอนุมานว่าผู้ชมมีแนวโน้มที่จะไม่ชอบร้านอาหารที่กำลังค้นหาอยู่

ปัจจุบันจึงยังทำการค้นหาร้านอาหารต่อไป โดยให้คะแนนความชื่นชอบเท่ากับ -0.5 จากนั้นระบบจะทำการรวมคะแนนของผู้ใช้แต่ละคน

ข้อจำกัด: ไม่สามารถนำมาประยุกต์ใช้กับองค์กรที่ยังไม่มีการเก็บข้อมูล Activity ของผู้ใช้ในหน้าเว็บไซต์ เหมาะสำหรับองค์กรที่มีระบบออนไลน์เท่านั้น และเนื่องด้วยข้อจำกัดของข้อมูลที่น่าวิเคราะห์ยังไม่มีเก็บข้อมูลพฤติกรรมของลูกค้านี่ออนไลน์และเว็บไซต์ร้านค้า ผู้วิจัยจึงนำเสนอวิธีอื่นที่เหมาะสมกับข้อมูลที่มีอยู่ของธุรกิจ

2.2.2 บริบท (Context)

ในงานวิจัยเกี่ยวกับระบบแนะนำส่วนใหญ่มักจะไม่ค่อยให้ความสำคัญกับ “บริบท” (Context) (Anand & Mobasher, 2007) โดยระบบแนะนำส่วนมากจะให้ความสนใจเพียง 2 สิ่งคือ สิ่งของที่จะแนะนำ (Items) และผู้ใช้ (Users) โดยไม่ได้คำนึงถึงบริบทแวดล้อมเข้ามาเกี่ยวข้อง ยกตัวอย่างเช่น เวลา หรือ สถานที่ แต่ในความเป็นจริงนั้นผู้ใช้นั้นจะมีการตอบสนองหรือการให้คะแนนความชอบ (Rating) ในบริบทที่ต่างกันแตกต่างกันออกไป Lieberman and Selker (2000) ได้นิยามเกี่ยวกับ “บริบท” ไว้ว่าบริบทคือทุกสิ่งที่มีผลต่อการแนะนำของระบบ ยกตัวอย่างเช่น ระบบแนะนำสถานที่ท่องเที่ยวควรจะแนะนำสถานที่ท่องเที่ยวที่ต่างกันระหว่างฤดูร้อนกับฤดูหนาว หรือระบบแนะนำร้านอาหารควรจะแนะนำร้านอาหารสำหรับลูกค้าที่มาเป็นกลุ่มกับลูกค้าที่มาคนเดียวต่างกัน

สอดคล้องกับงานวิจัยทางการตลาด ซึ่งทำการวิจัยด้านพฤติกรรมของผู้บริโภคเกี่ยวกับการตัดสินใจเชื่อว่าผู้บริโภคมีการตัดสินใจที่ไม่แน่นอนขึ้นอยู่กับบริบทแวดล้อม หรือกล่าวได้ว่าผู้บริโภคคนเดิมอาจจะตัดสินใจเปลี่ยนไปและมีความชอบในสินค้าหรือแบรนด์ที่ต่างออกไปในบริบทที่ต่างกัน (Klein & Yadav, 1989; Lussier & Olshavsky, 1979; Robertson & Kassarian, 1991) เช่นเดียวกับ Lilien (1992) ได้กล่าวว่าผู้บริโภคมีการตัดสินใจที่ไม่แน่นอน ขึ้นอยู่กับสถานการณ์ในขณะนั้น สินค้าหรือบริการ เช่นการตัดสินใจซื้อสินค้าสำหรับตนเอง สินค้าสำหรับครอบครัว หรือสินค้าที่ใช้เป็นของขวัญ รวมถึงสถานการณ์ที่ซื้อของ เช่นการซื้อสินค้าจากสมุด Catalog หรือการซื้อสินค้าโดยมีพนักงานขายให้คำแนะนำ ดังนั้นความถูกต้องของการทำนายความชอบของผู้บริโภคจะขึ้นอยู่กับระดับความเกี่ยวข้องของบริบทที่นำมาใช้ในระบบแนะนำสินค้า

โดยตัวอย่างขององค์กรที่มีการนำบริบทมาใช้ร่วมกับระบบแนะนำ เช่น ระบบแนะนำรายการของ Netflix ที่ไม่เพียงแต่ใช้ข้อมูลพฤติกรรมของผู้ใช้เท่านั้น แต่ได้มีการนำบริบทเข้ามาใช้ร่วมด้วย เช่น ช่วงเวลา ความแปลกใหม่ และความหลากหลาย ซึ่งเป็นปัจจัยสำคัญที่ทำให้ Netflix มีข้อมูลจำนวนมหาศาลเพื่อใช้ในการหารูปแบบการใช้งานของผู้ใช้ เพื่อนำมาใช้แนะนำรายการให้แก่ผู้ใช้ได้ดียิ่งขึ้น

การนำบริบทมาใช้ร่วมกับระบบแนะนำ (Context-aware Recommendation) สามารถแบ่งออกเป็น 3 วิธี ได้แก่ Pre-filtering, Post-filtering และ Contextual Modeling โดยวิธี Pre-filtering จะทำการเลือกข้อมูลที่อยู่ในบริบทที่กำหนดออกมาก่อนแล้วจึงทำการแนะนำสินค้า ส่วนวิธี Post-filtering จะทำการแนะนำสินค้าจากระบบแนะนำสินค้านั้นก่อนตามวิธีการปกติ จากนั้นจึงปรับค่าสำหรับผู้ใช้งานแต่ละคนโดยใช้ข้อมูลตามบริบทที่ต้องการ แต่หากเปรียบเทียบกับสองวิธีแรก วิธี Contextual Modeling จะเป็นการใช้ข้อมูลบริบทในกระบวนการสร้าง Model โดยตรง (Oku, Nakajima, Miyazaki, & Uemura, 2006) โดยงานวิจัยของ Panniello, Tuzhilin, Gorgoglione, Palmisano, and Pedone (2009) ได้ทำการเปรียบเทียบประสิทธิภาพระหว่างวิธี Pre-filtering และ Post-filtering ซึ่งพบว่าไม่มีวิธีใดที่ดีกว่ากันอย่างชัดเจน

2.2.3 ปัญหาผู้ใช้งานระบบใหม่ (Cold-start)

แบ่งเป็น 3 หัวข้อในการทบทวนวรรณกรรม ได้แก่ การใช้ข้อมูลด้านบุคลิกภาพของผู้ใช้ระบบ การใช้ข้อมูลจากสื่อสังคมออนไลน์ (Social Media) ของผู้ใช้งาน และการใช้ข้อมูลประชากรศาสตร์ของผู้ใช้ระบบมาช่วยในการแก้ปัญหาผู้ใช้งานระบบรายใหม่

2.2.3.1 การใช้ข้อมูลด้านบุคลิกภาพของผู้ใช้ระบบ (Personality-based Recommendation System)

เป็นวิธีการที่เพิ่มคุณลักษณะ (Attribute) ด้านบุคลิกภาพของลูกค้าย (Personality Traits) ลงในข้อมูลส่วนตัว (Profile) ของลูกค้ายรายใหม่ (Braunhofer, Elahi, & Ricci, 2015) โดยมีหลักฐานที่แสดงว่าคนที่มีความคล้ายคลึงกันจะมีกระบวนการตัดสินใจที่ต่างกัน (Fernández-Tobías, Braunhofer, Elahi, Ricci, & Cantador, 2016) รวมถึงคนที่มีความคล้ายคลึงกันก็จะมีรสนิยมที่คล้ายกัน (Schedl, Zamani, Chen, Deldjoo, & Elahi, 2018) อีกทั้งงานวิจัยของ Rentfrow and Gosling (2003) ได้ทำการศึกษาพบว่าบุคลิกภาพมีความสัมพันธ์กับเพลงที่ชื่นชอบ ยกตัวอย่างเช่น คนที่มีลักษณะบุคลิกภาพแบบเปิดเผยจะชอบฟังเพลงประเภท Jazz Blue และ Classic ส่วนคนที่มีบุคลิกภาพแบบชอบเข้าสังคมและมีความเป็นมิตรจะชอบประเภทเพลง Rap หรือ Hip-hop ซึ่งต่อมา Hu and Pu (2011) ได้นำผลลัพธ์จากงานวิจัยของ Rentfrow and Gosling (2003) มาสร้างระบบแนะนำเพลงโดยใช้ข้อมูลลักษณะบุคลิกภาพของผู้ใช้มาช่วยในการแนะนำแนวเพลงที่ผู้ใช้งานรายใหม่จะชื่นชอบ

ข้อจำกัดสำหรับวิธี Personality-based คือระบบจะต้องมีข้อมูลด้านบุคลิกภาพของผู้ใช้ หรืออาจจะจะต้องให้ผู้ใช้งานรายใหม่ทำแบบสอบถามเกี่ยวกับบุคลิกภาพก่อน ซึ่งเป็นส่วนที่ต้องใช้เวลาและผู้ใช้งานอาจไม่ให้ความร่วมมือ

2.2.3.2 การใช้ข้อมูลจากสื่อสังคมออนไลน์ของผู้ใช้ระบบ (Social Media Approach)

เป็นวิธีการนำข้อมูลของผู้ใช้ระบบจาก Social Media มาช่วยแก้ปัญหาผู้ใช้งานระบบรายใหม่ (Cold-start) ยกตัวอย่างเช่นในงานวิจัยของ Nair, Moh, and Moh (2016) หาความสนใจของผู้ใช้รายใหม่จากการพิมพ์ข้อความใน Twitter เพื่อเปรียบเทียบความสนใจของผู้ใช้กับ Metadata ของหนังสือ (พล็อตหนังสือ) จากนั้นทำการหาความสัมพันธ์ระหว่างผู้ใช้รายใหม่กับประเภทของหนังสือ ซึ่งคล้ายกับงานวิจัยของ Rana, Salima, Usama, and Hammam (2014) ที่ทำการสกัด (Extract) หาความสนใจของผู้ใช้จากการพิมพ์ข้อความใน Twitter จากนั้นนำมาเปรียบเทียบหาความคล้ายคลึงระหว่างความสนใจของผู้ใช้รายใหม่กับผู้ใช้รายเก่าในระบบ และทำการแนะนำหนังสือที่คาดว่าผู้ใช้รายใหม่จะชื่นชอบ

ข้อจำกัดสำหรับวิธี Social Media คือเรื่องความเป็นส่วนตัวของข้อมูลโดยหากลูกค้าไม่ยินยอมที่จะเปิดเผยข้อมูลก็จะทำให้องค์กรไม่สามารถเข้าถึงข้อมูลส่วนตัวของผู้ใช้ได้

2.2.3.3 การใช้ข้อมูลประชากรศาสตร์ของผู้ใช้ระบบ (Demographic-based)

วิธี Demographic-based Recommendation System เป็นวิธีที่แนะนำสินค้าโดยใช้ข้อมูลทางประชากร (Pazzani, 1999) ซึ่งเป็นวิธีที่ไม่จำเป็นต้องมีข้อมูลคะแนนความชื่นชอบ (Rating) จากผู้ใช้รายใหม่แต่จะแนะนำสินค้าโดยอ้างอิงจากข้อมูลประชากรศาสตร์แทน โดยมีข้อสมมติฐานว่าผู้ใช้ที่มีข้อมูลลักษณะทางประชากรศาสตร์ที่คล้ายกันจะมีความชื่นชอบและรสนิยมที่คล้ายกัน (J. Bobadilla, Ortega, Hernando, & Gutiérrez, 2013) ซึ่งสอดคล้องกับงานวิจัยของ Bhatia (2016) และ Pandey and Rajpoot (2016) ที่สามารถนำข้อมูลทางประชากรศาสตร์ของผู้ใช้ระบบรายใหม่มาใช้ในการแก้ปัญหาผู้ใช้งานระบบรายใหม่ได้

2.2.4 โครงข่ายประสาทเทียม (Artificial Neural Network)

ปัจจุบันการวิจัยของระบบแนะนำมุ่งไปที่แนวทางการเรียนรู้เชิงลึกมากขึ้นเนื่องจากมีหลักฐานที่แสดงว่าวิธี Artificial Neural Network (ANN) มีประสิทธิภาพในการทำงานของระบบแนะนำที่ดีกว่าวิธี Collaborative Filtering (Jesus Bobadilla, Alonso, & Hernando, 2020; He et al., 2017) ซึ่งวิธี ANN ได้ถูกนำไปใช้กับลักษณะของข้อมูลที่หลากหลาย เช่น รูปภาพ ข้อความ เพลง หรือวิดีโอ ตัวอย่างเช่น Oord, Dieleman, and Schrauwen (2013) ใช้ประโยชน์จากโครงข่ายประสาทเทียมแบบ Deep Convolutional เพื่อใช้แนะนำเพลง หรือการใช้ Deep Neural Network เพื่อแนะนำ Videos บน Youtube (Covington et al., 2016) ส่วนงานวิจัยของ Chu and Tsai (2017) แนะนำร้านอาหารให้แก่ลูกค้าโดยใช้ข้อมูลรูปภาพที่ลูกค้าใช้รีวิวร้านอาหารบน Blogs

บทที่ 3

วิธีการดำเนินการวิจัย

ในส่วนที่ 3 ผู้วิจัยนำเสนอวิธีการดำเนินการวิจัย โดยจะแบ่งออกเป็น 2 ส่วนหลักได้แก่การเก็บรวบรวมและจัดเตรียมข้อมูล (Dataset) และระบบแนะนำสินค้าแบบผสม (A Hybrid Recommendation System) โดยจะแบ่งออกเป็น 2 ส่วนคือ 1) ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่ และ 2) ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า โดยจะกล่าวถึงรายละเอียดของทั้งสองส่วนในลำดับต่อไป

3.1 การเก็บรวบรวมและจัดเตรียมข้อมูล (Dataset)

ชุดข้อมูลที่นำมาวิเคราะห์ในงานวิจัยนี้ผู้วิจัยได้รับชุดข้อมูลจากธุรกิจร้านอาหารหนึ่งในประเทศไทย โดยใช้ข้อมูลธุรกรรม (Transactions) ตั้งแต่เดือนมกราคมถึงเดือนธันวาคม พ.ศ. 2562 โดยนำมาวิเคราะห์เฉพาะกรณีการสั่งซื้ออาหารประเภทจานหลักของลูกค้าที่เป็นสมาชิกร้านอาหารเท่านั้น โดยข้อมูลที่นำมาวิเคราะห์มีจำนวนข้อมูลเท่ากับ 344,373 ธุรกรรม (Transactions) ส่วนข้อมูลส่วนตัวของลูกค้าได้จากการเก็บข้อมูลผ่านบัตรสมาชิกของร้านอาหารซึ่งข้อมูลบัตรสมาชิกที่นำมาวิเคราะห์มีจำนวนทั้งสิ้น 66,248 ใบ และเนื่องจากเป็นชุดข้อมูลจริงจึงทำให้เกิดข้อจำกัดในบางประการ เช่น การเข้าถึงข้อมูลของลูกค้าและแหล่งข้อมูลที่ต้องการมีอยู่ในปัจจุบันรวมถึงลักษณะของข้อมูลที่มีจำนวนว่างค่อนข้างมาก (Sparse Data) ซึ่งเป็นส่วนสำคัญที่อาจส่งผลต่อการวิเคราะห์และค่าความถูกต้องของระบบแนะนำ

3.1.1 การจัดเตรียมข้อมูล (Data Preprocessing)

ในส่วนข้อมูลส่วนตัวของลูกค้าที่ได้จากการเก็บข้อมูลผ่านบัตรสมาชิกของร้านอาหาร ผู้วิจัยได้จัดทำ Data Cleaning ดังนี้

- 1) ตัวแปรเพศ เนื่องจากข้อมูลที่ได้จากบัตรสมาชิกซึ่งได้มีการกรอกข้อมูลในช่องระบุเพศนั้น ลูกค้ามีการตอบกลับค่อนข้างน้อย ผู้วิจัยจึงใช้ข้อมูลในส่วนของคำนำหน้าชื่อ (Title) มา

ช่วยในการแบ่งเพศของลูกค้ายกนี้ ลูกค้าที่มีค่านำหน้าเด็กหญิง นางสาว และนางจะถูกกำหนดเป็นเพศหญิง ส่วนลูกค้าที่มีค่านำหน้าเด็กชาย และนายจะถูกกำหนดเป็นเพศชาย

2) ตัวแปรอายุ เนื่องจากการเก็บข้อมูลจากบัตรสมาชิกจะไม่มีช่องที่ให้ลูกค้าระบุอายุอย่างชัดเจน แต่จะเก็บข้อมูลวันเดือนปีเกิดของผู้ถือบัตร ดังนั้นผู้วิจัยจึงทำการดึงค่าปีเกิดที่ลูกค้าตอบกลับมาใช้ ซึ่งจะสามารถแบ่งได้เป็น 2 กรณี คือกรณีที่ลูกค้าใส่ปีเกิดแบบ พ.ศ. และกรณีที่ลูกค้าใส่ปีเกิดแบบ ค.ศ. จากนั้นจะนำมาแปลงค่าเป็นอายุของลูกค้าโดยเทียบกับปี พ.ศ. 2562 หรือ ค.ศ. 2019 ซึ่งเป็นปีที่น่าข้อมูลมาวิเคราะห์ ภายหลังจากที่ได้อายุของลูกค้าแต่ละคนแล้ว ผู้วิจัยจะแบ่งกลุ่มอายุออกเป็น Generation ทั้งหมด 4 กลุ่มได้แก่ Baby Boomer, Gen X, Gen Y และ Gen Z โดยในแต่ละกลุ่มมีช่วงอายุดังนี้

- (1) กลุ่ม Baby Boomer คือคนที่เกิดในช่วงปี พ.ศ. 2489 – 2507
- (2) กลุ่ม Gen X คือคนที่เกิดในช่วงปี พ.ศ. 2508 – 2522
- (3) กลุ่ม Gen Y คือคนที่เกิดในช่วงปี พ.ศ. 2523 – 2540
- (4) กลุ่ม Gen Z คือคนที่เกิดในช่วงปี พ.ศ. 2541 – 2552

โดยผู้วิจัยจะตัดกลุ่มค่าส่วนเกิน (Outlier) ในกรณีที่เจ้าของบัตรมีอายุต่ำกว่า 10 ปี เนื่องจากในกลุ่มคนที่มีอายุน้อยกว่า 10 ปีอาจเป็นผู้ปกครองที่สมัครบัตรในชื่อของลูก ซึ่งอาจก่อให้เกิด Bias หากนำมาวิเคราะห์ร่วมด้วย

3) ตัวแปรอาชีพ แบ่งออกเป็น 5 กลุ่มได้แก่ ลูกจ้าง, ธุรกิจส่วนตัว, เจ้าหน้าที่รัฐ, นักเรียน และอื่น ๆ โดยเป็นตัวเลือกที่ทางร้านอาหารมีให้ลูกค้าเลือกกรอกในขั้นตอนการสมัครบัตรสมาชิก ดังนั้นในขั้นตอนของการ Clean ข้อมูลจึงไม่ต้องปรับเปลี่ยนในส่วนของตัวแปรอาชีพ

4) ตัวแปรภูมิลำเนา ในการเก็บข้อมูลจากบัตรสมาชิกผู้สมัครบัตรจะต้องกรอกข้อมูลในส่วนของที่อยู่อาศัย โดยผู้วิจัยได้นำข้อมูลเฉพาะในส่วนจังหวัดที่ลูกค้าอาศัยอยู่มาใช้ จากนั้นจะแปลงจากจังหวัดให้กลายเป็นภูมิภาค โดยทำการ Map จังหวัดที่ลูกค้าอาศัยอยู่เข้ากับภูมิภาคของจังหวัดนั้น ๆ โดยจังหวัดกรุงเทพฯ นนทบุรี สมุทรปราการ สมุทรสาคร นครปฐม และปทุมธานี จะถูกจัดให้อยู่กลุ่มเดียวกันคือ กรุงเทพฯและปริมณฑล ส่วนจังหวัดอื่น ๆ จะจัดตามภูมิภาคปกติ

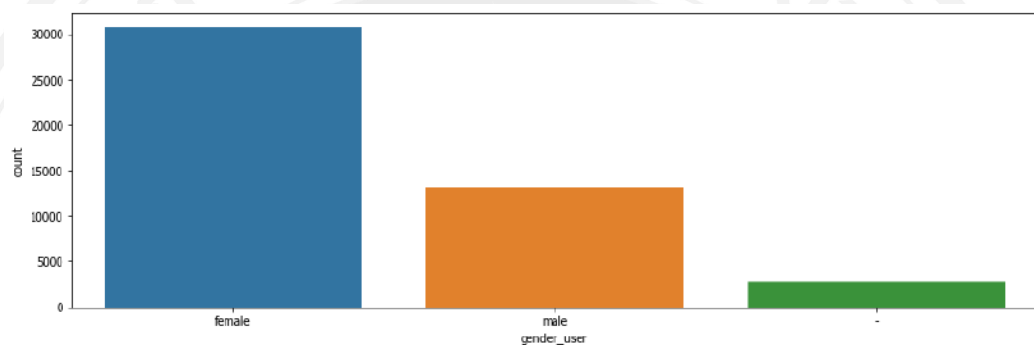
ตารางที่ 3.1 ตัวอย่างข้อมูลนำเข้าหลังจากทำ Data Cleaning

Member Id	เพศ	อายุ	อาชีพ	ภูมิลำเนา
001	หญิง	Gen X	ลูกจ้าง	ก.ท.ม.
002	หญิง	Gen Y	ธุรกิจส่วนตัว	ภาคกลาง

Member Id	เพศ	อายุ	อาชีพ	ภูมิลำเนา
003	ชาย	Gen Z	นักเรียน	ภาคเหนือ

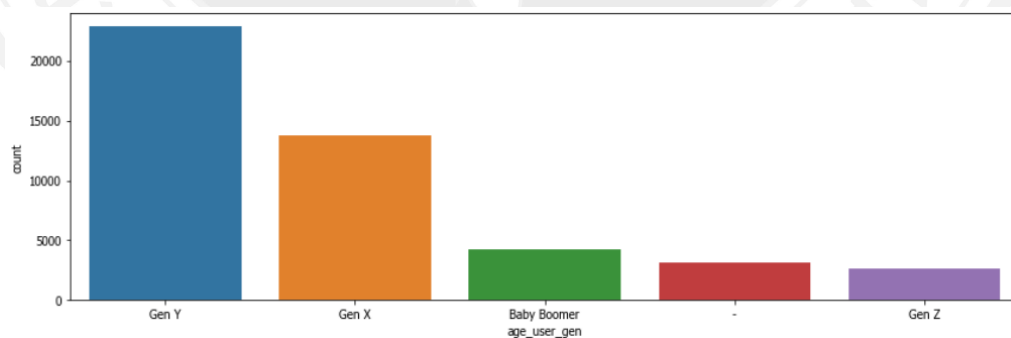
3.1.2 การแสดงข้อมูล (Data Visualization)

ภายหลังจากทำ Data Preprocessing เรียบร้อยแล้ว พบว่าตัวแปรประชากรแต่ละตัวมีการกระจายตัวดังต่อไปนี้



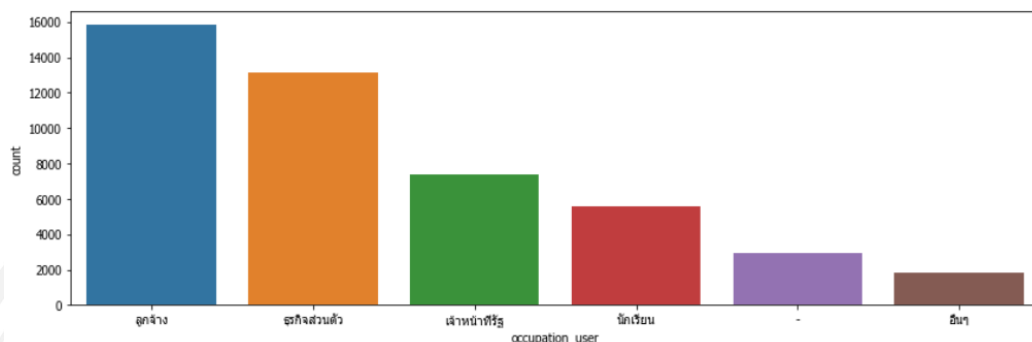
ภาพที่ 3.1 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรเพศ

ภาพที่ 3.1 แสดงการกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรเพศ ซึ่งพบว่าลูกค้าที่เป็นสมาชิกของร้านอาหารส่วนใหญ่เป็นลูกค้าเพศหญิง โดยมีจำนวนมากกว่าเพศชายเกือบ 1 เท่าตัว



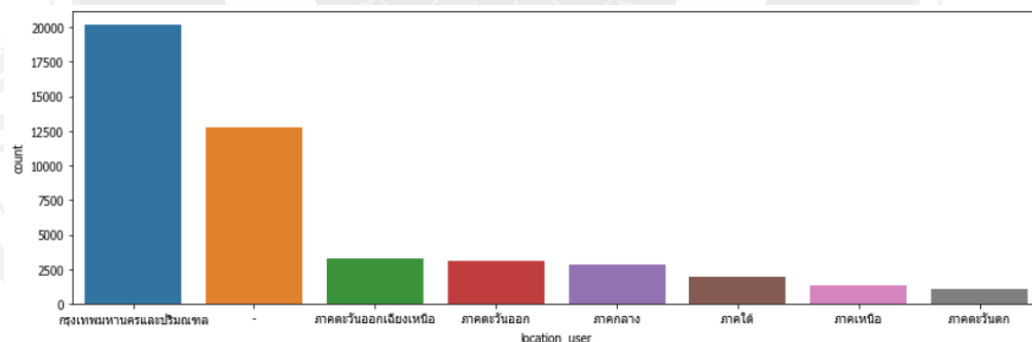
ภาพที่ 3.2 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรอายุ

ภาพที่ 3.2 แสดงการกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรอายุ ซึ่งกลุ่มลูกค้าที่มีจำนวนมากที่สุด คือ ลูกค้าที่มีอายุในช่วง Generation Y รองลงมาคือกลุ่ม Generation X ส่วนกลุ่ม Baby Boomer และ Generation Z มีจำนวนค่อนข้างน้อยเมื่อเทียบกับสองกลุ่มแรก



ภาพที่ 3.3 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรอาชีพ

ภาพที่ 3.3 แสดงการกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรอาชีพ ซึ่งลูกค้าที่เป็นสมาชิกส่วนใหญ่ของร้านอาหารประกอบอาชีพลูกจ้าง รองลงมาคือประกอบธุรกิจส่วนตัว เจ้าหน้าที่รัฐ นักเรียน และอื่น ๆ ตามลำดับ



ภาพที่ 3.4 การกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรภูมิภาค

ภาพที่ 3.4 แสดงการกระจายตัวของจำนวนสมาชิกร้านอาหารผ่านตัวแปรภูมิภาค พบว่าลูกค้าที่เป็นสมาชิกของร้านอาหารอาศัยอยู่ในเขตกรุงเทพและปริมณฑลเป็นส่วนมาก ส่วนในภูมิภาคอื่น ๆ มีจำนวนเฉลี่ยใกล้เคียงกันและค่อนข้างน้อยเมื่อเทียบกับกรุงเทพและปริมณฑล โดยภาคตะวันตกเป็นภูมิภาคที่มีลูกค้าสมาชิกร้านอาหารน้อยที่สุด

3.2 ระบบแนะนำสินค้าแบบผสม (A Hybrid Recommendation System)

3.2.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่

จากปัญหาผู้ใช้งานระบบรายใหม่ (Cold-start Problem) แม้จะมีงานวิจัยพยายามแก้ปัญหาด้วยวิธีที่แตกต่างกันออกไป แต่วิธีเหล่านั้นก็ยังพบข้อจำกัดและความเหมาะสมของการนำไปใช้กับแต่ละองค์กรที่ต่างกัน ดังนั้นในงานวิจัยนี้ผู้วิจัยจึงเลือกใช้วิธีที่เหมาะสมกับข้อมูลที่ต้องการมีอยู่อย่างจำกัดมากที่สุด โดยผู้วิจัยนำเสนอวิธีการแก้ปัญหาโดยใช้วิธี Demographic-based Recommendation System ซึ่งเป็นวิธีที่สามารถใช้แก้ปัญหากรณีลูกค้ารายใหม่ได้ เนื่องจากไม่จำเป็นต้องใช้ข้อมูลคะแนนความชื่นชอบ (Rating) จากผู้ใช้รายใหม่ แต่จะแนะนำสินค้าโดยอ้างอิงจากข้อมูลประชากรศาสตร์แทน (Pazzani, 1999) โดยมีข้อสมมติฐานว่าผู้ใช้ที่มีข้อมูลลักษณะทางประชากรศาสตร์ที่คล้ายกันจะมีความชื่นชอบและรสนิยมที่คล้ายกัน (J. Bobadilla et al., 2013) นอกจากนั้นจะใช้ร่วมกับฟังก์ชันอายุ (Similarity Age Function) และวิธีการเรียนรู้ของเครื่องจักร (Machine Learning) ได้แก่วิธี Clustering และ Classification รวมถึงการนำบริบท (Contexts) มาใช้ร่วมกับโมเดล โดยบริบทที่ผู้วิจัยเลือกใช้แบ่งออกเป็น 3 โดเมนได้แก่ 1) บริบทด้านเวลา (Temporal) (Bisogni et al., 2007) 2) บริบทด้านภูมิศาสตร์ (Geographic) และ 3) บริบทด้านร้านค้า (Store Profiles) โดยในแต่ละด้านจะทำการเลือกตัวแปรที่เกี่ยวข้องกับด้านนั้น ๆ เพื่อนำมาใช้เป็นบริบทที่ต่างกันออกไป

1) บริบทด้านเวลา (Temporal)

ตัวแปรที่นำมาพิจารณาในบริบทด้านเวลาประกอบด้วย 2 ตัวแปรคือ ช่วงเวลาของวัน (Time of the Day) และประเภทของวัน (Weekday/Weekend) โดยตัวแปรช่วงเวลาของวันจะถูกแบ่งออกตามการกระจายตัวของการมาทานอาหารของลูกค้า โดยจะแบ่งเป็น 4 ช่วงเวลาดังนี้

ช่วงเวลาในกลุ่มที่ 1 คือช่วงเวลาก่อน 10.59 น.

ช่วงเวลาในกลุ่มที่ 2 คือช่วงเวลา 11.00 – 13.59 น.

ช่วงเวลาในกลุ่มที่ 3 คือช่วงเวลา 14.00 – 16.59 น.

ช่วงเวลาในกลุ่มที่ 4 คือช่วงเวลา 17.00 น. เป็นต้นไป

ส่วนตัวแปรประเภทของวันจะถูกแบ่งออกเป็น 2 กลุ่มดังนี้

กลุ่มวันธรรมดา ประกอบด้วยวันจันทร์ถึงวันศุกร์ และ

กลุ่มวันหยุด ประกอบด้วยวันเสาร์และวันอาทิตย์

2) บริบทด้านภูมิศาสตร์ (Geographic)

ตัวแปรที่นำมาพิจารณาในบริบทด้านภูมิศาสตร์ประกอบด้วย 2 ตัวแปรคือ ประเภทห้างที่ร้านค้าตั้งอยู่ และภูมิภาคที่ตั้งของร้าน โดยตัวแปรประเภทห้างที่ร้านค้าตั้งอยู่จะถูกแบ่งออกเป็น 4 กลุ่มดังนี้

- กลุ่มที่ 1 ร้านค้าตั้งอยู่ในห้างประเภท Community Mall
- กลุ่มที่ 2 ร้านค้าตั้งอยู่ในห้างประเภท Department Store
- กลุ่มที่ 3 ร้านค้าตั้งอยู่ในห้างประเภท Hypermarket
- กลุ่มที่ 4 ร้านค้าตั้งอยู่ในห้างประเภท อื่น ๆ

เนื่องจากในบริบทของสถานที่ที่ต่างกัน พฤติกรรมการทานอาหารหรือรายการอาหารที่ขายก็อาจมีความแตกต่างกัน ยกตัวอย่างเช่น สาขาที่ตั้งอยู่ที่โรงพยาบาลหรือสนามบินอาจมีช่วงเวลา Peak Hour ที่แตกต่างจากสาขาที่ตั้งอยู่ในห้างสรรพสินค้า หรือแม้กระทั่งสาขาที่ตั้งอยู่ในห้างสรรพสินค้าเหมือนกันก็ยังสามารถแบ่งออกตามประเภทของห้างสรรพสินค้าได้อีก โดยห้างสรรพสินค้าประเภทห้าง Hypermarket กับห้างสรรพสินค้าประเภทห้าง Department Store ก็จะมีลักษณะของกลุ่มลูกค้าที่แตกต่างกัน

ส่วนตัวแปรภูมิภาคที่ตั้งของร้านจะถูกแบ่งออกเป็น 2 กลุ่มดังนี้

กลุ่มที่ 1 สาขากรุงเทพและปริมณฑล ประกอบด้วยจังหวัดกรุงเทพฯ นนทบุรี สมุทรปราการ สมุทรสาคร นครปฐม และปทุมธานี

กลุ่มที่ 2 สาขาต่างจังหวัด ประกอบด้วยจังหวัดอื่น ๆ ที่เหลือทั้งหมด

3) บริบทด้านร้านค้า (Store Profiles)

ตัวแปรที่นำมาพิจารณาในบริบทด้านข้อมูลร้านค้าประกอบด้วย 2 ตัวแปรคือ เกรดของร้านค้า และประเภทการสั่งอาหาร โดยตัวแปรเกรดของร้านค้าจะถูกแบ่งออกตามยอดขายรวมของสาขานั้น ๆ ซึ่งจะถูกรวบรวมออกเป็น 6 กลุ่มดังนี้

- กลุ่มที่ 1 คือร้านค้าที่มียอดขายของสาขารวมในระดับ AAA
- กลุ่มที่ 2 คือร้านค้าที่มียอดขายของสาขารวมในระดับ AA
- กลุ่มที่ 3 คือร้านค้าที่มียอดขายของสาขารวมในระดับ A
- กลุ่มที่ 4 คือร้านค้าที่มียอดขายของสาขารวมในระดับ B
- กลุ่มที่ 5 คือร้านค้าที่มียอดขายของสาขารวมในระดับ C
- กลุ่มที่ 6 คือร้านค้าที่มียอดขายของสาขารวมในระดับ D

ส่วนตัวแปรประเภทการสั่งอาหารจะถูกแบ่งออกเป็น 2 กลุ่มดังนี้

- กลุ่มที่ 1 คือกรณีลูกค้าทานที่ร้าน (Dine-in)
- กลุ่มที่ 2 คือกรณีลูกค้าสั่งอาหารกลับบ้าน (Take-away)

โดยวิธีการแนะนำสินค้าสำหรับกรณีลูกค้ารายใหม่จะแบ่งออกเป็น 2 ขั้นตอนหลักดังนี้

3.2.1.1 แบ่งกลุ่มลูกค้ารายเก่า

ขั้นตอนแรกจะทำการแบ่งกลุ่มลูกค้ารายเก่าที่มีอยู่ในระบบตามลักษณะทางประชากรของลูกค้า เพื่อทำการแบ่งกลุ่มฐานลูกค้ารายเก่าของร้านอาหารและดูลักษณะทางประชากรของลูกค้าในแต่ละกลุ่ม ซึ่งตัวแปรที่นำมาใช้ในการแบ่งกลุ่มทุกตัวเป็นตัวแปรประเภทกลุ่ม (Category) ซึ่งได้ข้อมูลมาจากบัตรสมาชิกของลูกค้าร้านอาหาร ดังนั้นผู้วิจัยจึงแบ่งกลุ่มของลูกค้าเก่าในระบบด้วยวิธี K-modes Clustering โดยจะทำการเลือกจำนวนกลุ่มหรือ Cluster จากวิธี Elbow Method โดยตัวแปรด้านประชากรศาสตร์ที่นำมาใช้ในการแบ่งกลุ่มทั้งหมด 4 ตัวแปรแต่ละตัวแปรมีค่าดังนี้

ตัวแปรเพศ แบ่งออกเป็น 3 กลุ่ม ได้แก่ เพศชาย, เพศหญิง และไม่ระบุ

ตัวแปรอายุ แบ่งออกเป็น 5 กลุ่ม ได้แก่ Baby Boomer, Gen X, Gen Y, Gen Z และไม่ระบุ

ตัวแปรอาชีพ แบ่งออกเป็น 6 กลุ่ม ได้แก่ ลูกจ้าง, ธุรกิจส่วนตัว, เจ้าหน้าที่รัฐ, นักเรียน, อื่น ๆ และไม่ระบุ

ตัวแปรภูมิลำเนา แบ่งออกเป็น 8 กลุ่ม ได้แก่ กรุงเทพมหานคร, ภาคกลาง, ภาคตะวันออก, ภาคตะวันตก, ภาคใต้, ภาคเหนือ, ภาคตะวันออกเฉียงเหนือ และไม่ระบุ

3.2.1.2 แนะนำสินค้าสำหรับลูกค้ารายใหม่

หลังจากทำการแบ่งกลุ่มลูกค้ารายเก่าในระบบตามลักษณะทางประชากรเรียบร้อยแล้ว ขั้นตอนถัดมาจะเป็นการแนะนำสินค้าสำหรับลูกค้ารายใหม่ โดยจะมีขั้นตอนในการแนะนำสินค้าทั้งหมด 4 ขั้นตอนดังต่อไปนี้

- ขั้นที่หนึ่ง ระบบจะใช้โมเดลต้นไม้หรือ Decision Tree Model เพื่อทำนายกลุ่มของลูกค้ารายใหม่ที่เข้าสู่ระบบว่ามีลักษณะทางประชากรคล้ายกับลูกค้าเก่าในกลุ่มใดตามที่ได้ถูกแบ่งไว้ในขั้นตอนก่อนหน้านี้มากที่สุด โดยขั้นตอนในการสร้างโมเดลจะใช้ข้อมูลตัวแปรประชากรศาสตร์ 4 ตัวแปร ได้แก่ เพศ อายุ อาชีพ และภูมิลำเนา ซึ่งผู้วิจัยใช้ข้อมูลจากบัตรสมาชิกของลูกค้า โดยบัตรสมาชิกที่นำมาวิเคราะห์มีจำนวนทั้งสิ้น 66,248 ใบ ซึ่งจะแบ่งออกมาครึ่งหนึ่งเพื่อใช้ในการสร้างโมเดลต้นไม้ประมาณ 33,000 ใบและอีกครึ่งหนึ่งใช้สำหรับในการทดสอบทั้งโมเดลแนะนำสินค้าให้แก่ลูกค้ารายใหม่ โดยสำหรับโมเดลต้นไม้ผู้วิจัยทำการแบ่งสัดส่วนข้อมูลบัตรสมาชิกของลูกค้าจาก 33,000 ใบด้วยสัดส่วน 70 ต่อ 30 เพื่อใช้เป็นชุดข้อมูลที่ใช้สอนโมเดล (Training Set) ร้อยละ 70 และชุดข้อมูลที่ใช้ทดสอบโมเดล (Test Set) อีกร้อยละ 30

ตารางที่ 3.2 ตัวอย่างตารางการทำนายกลุ่มให้แก่ลูกค้ารายใหม่ด้วยโมเดลต้นไม้

Member Id	เพศ	อายุ	อาชีพ	ภูมิลำเนา	กลุ่ม (y)
001	หญิง	Gen X	ลูกจ้าง	ก.ท.ม.	2
002	หญิง	Gen Y	ธุรกิจส่วนตัว	ภาคกลาง	1
003	ชาย	Gen Z	นักเรียน	ภาคเหนือ	7

• ขั้นที่สอง หลังจากโมเดลต้นไม้ทำนายกลุ่มของลูกค้ารายใหม่เรียบร้อยแล้ว ระบบจะใช้ Similarity Age Function ซึ่งเป็นฟังก์ชันในการกรองหาลูกค้าเก่าภายในกลุ่มที่มีอายุใกล้เคียงกับลูกค้ารายใหม่มากที่สุด โดย Similarity Age Function มีขั้นตอนการทำงาน 2 ขั้นตอนดังนี้

ขั้นที่ 1 เนื่องจากตัวแปรอายุเป็นตัวแปรประเภทกลุ่มโดยแบ่งเป็น Generation ดังนั้นในขั้นแรกระบบจะทำการกรองเฉพาะลูกค้ารายเก่าที่อยู่ใน Generation เดียวกันกับลูกค้ารายใหม่เท่านั้น

ขั้นที่ 2 เนื่องจากในแต่ละ Generation ก็ยังมีลูกค้าหลากหลายอายุ ดังนั้นในขั้นที่สองระบบจะกรองหาลูกค้ารายเก่าที่มีอายุเท่ากับหรือใกล้เคียงกับลูกค้ารายใหม่มากที่สุด

$$\min | \text{age}(\text{new}) - \text{age}(\text{old}) |$$

เมื่อ $\text{age}(\text{new})$ คือ อายุของลูกค้ารายใหม่ที่เข้าสู่ระบบ

$\text{age}(\text{old})$ คือ อายุของลูกค้ารายเก่าภายในกลุ่มเดียวกัน

• ขั้นที่สาม ระบบจะทำการกำหนดบริบทที่ตรงกับสถานการณ์ที่ต้องการแนะนำสินค้า โดยจะเลือกบริบทจากทั้ง 3 โดเมนได้แก่ โดเมนด้านเวลา ด้านภูมิศาสตร์ และด้านร้านค้า และจะนำข้อมูลในแต่ละด้านมาพิจารณาร่วมกัน (Cross Dimensions) เพื่อสร้างเป็นบริบทที่ต่างกันออกไป เป็นขั้นตอนสุดท้ายในการกรองข้อมูล (Post-filtering) โดยมีเงื่อนไขว่าหากไม่มีข้อมูลในบริบทที่กำหนดระบบจะข้ามไปที่ขั้นตอนต่อไป

• ขั้นตอนที่สี่ หลังจากทำการกรองข้อมูลครบทั้ง 3 ขั้นตอนแล้ว ระบบจะแนะนำรายการอาหารที่มีค่าเฉลี่ยของจำนวนครั้งในการสั่งสูงที่สุดจากจำนวนสมาชิกทั้งหมดที่เหลืออยู่

ตัวอย่างการแนะนำรายการอาหารกรณียลูกค้ารายใหม่ ระบบจะมีขั้นตอนและวิธีการในการแนะนำสินค้าดังนี้ จากขั้นตอนการแบ่งกลุ่มของลูกค้ารายเก่าของร้านอาหารสามารถแบ่งกลุ่มของลูกค้าออกได้เป็น 6 กลุ่ม โดยนาย A เป็นลูกค้ารายใหม่ของร้านอาหารที่มีลักษณะประชากรเป็นเพศ

ชาย มีอายุในช่วง Generation Y ประกอบอาชีพลูกจ้าง และอาศัยอยู่ในกรุงเทพมหานคร เมื่อเข้าสู่ระบบระบบจะใช้โมเดลต้นไม้ทำนายกลุ่มให้แก่ นาย A จากลักษณะทางประชากรโดยพบว่ามีลักษณะคล้ายคลึงกับลูกค้าเก่าในกลุ่มที่ 2 มากที่สุด จากนั้นภายในกลุ่มเดียวกันจะใช้ Similarity Age Function หาลูกค้ารายเก่าที่อายุใกล้เคียงกับนาย A มากที่สุด โดยจะเริ่มจากการกรองเฉพาะลูกค้าที่มีอายุอยู่ใน Generation Y จากนั้นภายในกลุ่มลูกค้า Generation Y ระบบจะหาลูกค้าที่อายุเท่ากับหรือใกล้เคียงกับนาย A มากที่สุด ในลำดับต่อไปจะใช้บริบทเป็นตัวกรองในขั้นตอนสุดท้าย เช่นหากนาย A มาทานอาหารที่ร้านในช่วงเวลาอาหารกลางวันในวันธรรมดา โดยเป็นสาขาที่มียอดขายของร้านประเภทเกรด A ซึ่งตั้งอยู่ในห้างสรรพสินค้าประเภท Hypermarket จังหวัดกรุงเทพมหานคร (Dine-in * Weekday * Lunch * Grade A * Hypermarket * Bangkok) ระบบก็จะกรองธุรกรรมการสั่งซื้อที่เกิดขึ้นเฉพาะในบริบทที่กำหนดไว้ สุดท้ายระบบจะแนะนำรายการอาหารที่มีค่าเฉลี่ยของจำนวนครั้งในการสั่งสูงที่สุดจากจำนวนสมาชิกทั้งหมดที่เหลืออยู่

ข้อมูลที่ใช้ในการทดสอบโมเดลแนะนำสินค้าสำหรับลูกค้ารายใหม่จะใช้ข้อมูลบัตรสมาชิกที่เหลืออยู่อีกครั้งหนึ่งจำนวน 33,000 ใบ โดยใช้สำหรับการทดสอบทั้งโมเดล ซึ่งผู้วิจัยได้จัดทำชุดข้อมูลที่เปรียบเสมือนการจำลองข้อมูลการสั่งซื้อจากลูกค้ารายใหม่ หรือ New User Test Set โดยการดึงข้อมูลการสั่งซื้อสินค้าเฉพาะการสั่งซื้อสินค้าครั้งแรกของแต่ละ Member ID เท่านั้น (First Purchase)

ผู้วิจัยจัดทำโมเดลแนะนำรายการอาหารสำหรับลูกค้ารายใหม่ออกเป็นสองกรณีดังนี้

- 1) โมเดลแนะนำรายการอาหารทั้ง 80 เมนู (ครบทุกรายการอาหารของร้านอาหาร)
- 2) โมเดลแนะนำรายการอาหารเฉพาะเมนูขายดี 10 เมนูแรกของร้าน

3.2.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า

ในกรณีที่องค์กรไม่มีการดำเนินการเก็บข้อมูลคะแนนความชื่นชอบโดยตรง (Explicit Rating) ในอดีตจากลูกค้ารายเก่าก็จะก่อให้เกิดปัญหาเช่นเดียวกัน เนื่องจากระบบไม่มีข้อมูลที่จะนำมาวิเคราะห์ความชื่นชอบของลูกค้า ดังนั้นในงานวิจัยนี้ผู้วิจัยได้นำเสนอวิธีการคิดคะแนนความชื่นชอบโดยนัย (Implicit Rating) ของลูกค้ารายเก่าด้วยวิธีที่เหมาะสมกับข้อมูลที่มีอยู่ในปัจจุบันของร้านอาหาร เพื่อใช้แทนคะแนนความชื่นชอบโดยตรงของลูกค้า และนำผลลัพธ์ที่ได้ไปใช้ในการสร้างระบบแนะนำสินค้าสำหรับลูกค้ารายเก่าในลำดับต่อไป ดังนั้นในส่วนของระบบแนะนำสินค้าสำหรับลูกค้ารายเก่าจะแบ่งออกเป็น 2 ส่วนหลักได้แก่ 1) ฟังก์ชันการคิดคะแนนความชื่นชอบโดยนัยของลูกค้ารายเก่า 2) ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า

3.2.2.1 ฟังก์ชันการคิดคะแนนความชื่นชอบโดยนัย (Implicit Rating Function) ของลูกค้ารายเก่า

ฟังก์ชันการคิดคะแนนคะแนนความชื่นชอบโดยนัยจะคำนวณโดยใช้ข้อมูลประวัติการสั่งซื้อสินค้าของลูกค้าแบบรายบุคคลดังนี้

$$\text{Implicit Rating Function (u,i)} = \frac{\text{จำนวนการสั่งเมนู i}}{\text{จำนวน transaction ทั้งหมดของผู้ใช้}} * 10$$

เมื่อ u คือ ผู้ใช้

i คือ รายการอาหาร

โดยค่าคะแนนความชื่นชอบโดยนัยที่ได้จะมีค่าอยู่ในช่วง 0 - 10 คะแนน ซึ่งสามารถแปลความหมายได้ดังนี้

0 คะแนน หมายถึงลูกค้าไม่มีความชื่นชอบในรายการอาหารนั้นเลย

10 คะแนน หมายถึงลูกค้ามีความชื่นชอบรายการอาหารนั้นมากที่สุด

3.2.2.2 แนะนำสินค้าสำหรับลูกค้ารายเก่า

ในส่วนของการขั้นตอนการสร้างระบบแนะนำสินค้าสำหรับลูกค้ารายเก่ามีขั้นตอนในการแนะนำสินค้าดังนี้

- ขั้นที่หนึ่ง ระบบจะทำการกำหนดบริบทที่ตรงกับสถานการณ์ที่ต้องการแนะนำสินค้า โดยจะเลือกบริบทจากทั้ง 3 โดเมนได้แก่ โดเมนด้านเวลา ด้านภูมิศาสตร์ และด้านร้านค้า และจะนำข้อมูลในแต่ละด้านมาพิจารณาร่วมกัน (Cross Dimensions) เพื่อสร้างเป็นบริบทที่ต่างกันไปเป็นขั้นตอนแรก (Pre-filtering) โดยมีเงื่อนไขว่าหากไม่มีข้อมูลในบริบทที่กำหนดระบบจะข้ามไปที่ขั้นตอนต่อไป

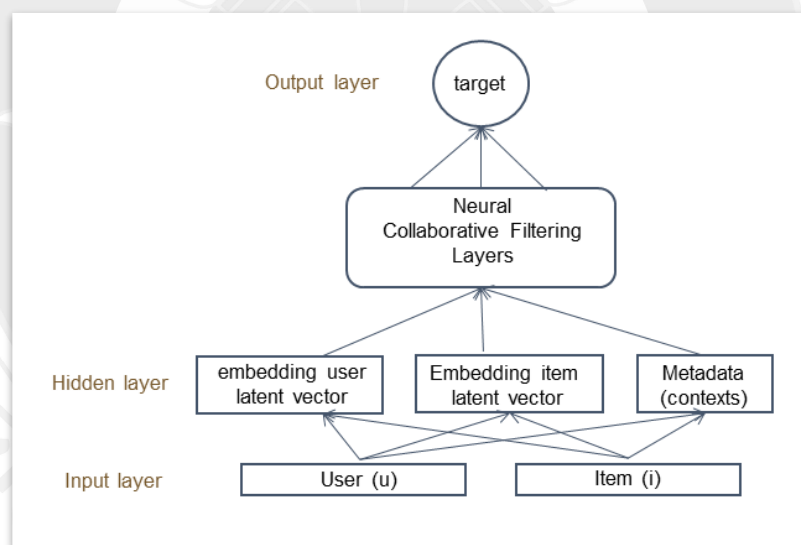
- ขั้นที่สอง เมื่อกำหนดบริบทที่ตรงกับสถานการณ์ที่ต้องการแนะนำสินค้าเรียบร้อยแล้ว ระบบจะทำการแปลงค่าจากประวัติการสั่งซื้อสินค้าของลูกค้าในบริบทที่กำหนดไว้ให้เป็นคะแนนความชื่นชอบโดยนัยของลูกค้าแต่ละบุคคลด้วยฟังก์ชันที่กำหนดไว้ในข้อ 3.2.2.1

- ขั้นที่สาม เมื่อมีคะแนนความชื่นชอบโดยนัยของลูกค้าแล้ว ระบบแนะนำสินค้าจะทำนายคะแนนความชื่นชอบของลูกค้าด้วยวิธี Collaborative Filtering (CF) แบบ Item-item Based

แต่เนื่องจากชุดข้อมูลที่นำมาวิเคราะห์เป็นชุดข้อมูลจริงจึงทำให้เกิดข้อจำกัดในเรื่องความว่างของข้อมูลที่มีจำนวนค่อนข้างมาก (Sparse Data) ผู้วิจัยจึงนำเสนอวิธี Artificial Neural Network (ANN) ร่วมกับการใช้บริบทในโมเดล (Contextual Modeling) อีกหนึ่งวิธีเพื่อใช้ในการทำนายคะแนนความชื่นชอบโดยนัยของลูกค้ารายเก่าร่วมด้วย เนื่องจากมีหลักฐานที่แสดงว่าวิธี Artificial

Neural Network มีประสิทธิภาพในการทำงานของระบบแนะนำที่ดีกว่าวิธี Collaborative Filtering ดังที่ผู้วิจัยทำการทบทวนงานวิจัยที่เกี่ยวข้องในบทที่ 2 ในส่วนของหัวข้อเรื่อง Artificial Neural Network สำหรับโมเดล Artificial Neural Network จะใช้ Adaptive Moment Estimation (Adam) เป็น Optimizer ของโมเดล และใช้ค่า MSE เป็น Loss Function ซึ่งจะทำให้การแปลงเป็นค่า RMSE ในภายหลังเพื่อใช้ในการเปรียบเทียบประสิทธิภาพการทำนายคะแนนความชื่นชอบของลูกค้ารายเก่าของโมเดลที่ได้จากทั้งสองวิธีและหาวิธีที่เหมาะสมกับลักษณะของข้อมูลมากที่สุด โดยโมเดล Artificial Neural Network มีสถาปัตยกรรมโครงสร้างของโมเดล (Architecture) อย่างคร่าวดังแสดงในภาพที่ 3.5

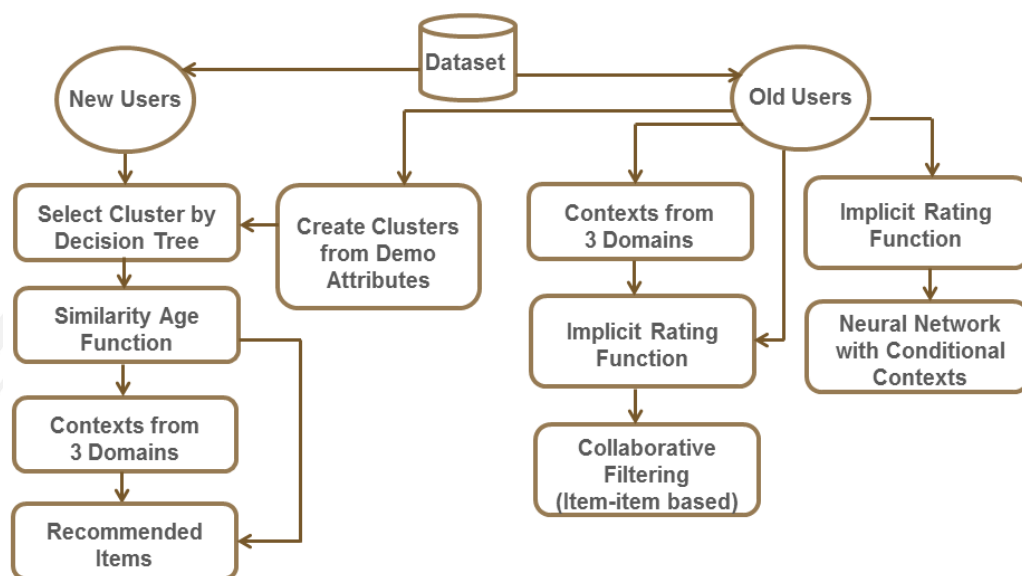
สำหรับชุดข้อมูลที่ใช้ในสร้างโมเดลผู้วิจัยใช้ข้อมูลธุรกรรมการสั่งซื้อสินค้าของลูกค้าที่เป็นสมาชิกของร้านอาหารในปี พ.ศ. 2562 ตั้งแต่เดือนมกราคม – เดือนธันวาคม โดยมีธุรกรรมการสั่งซื้อทั้งสิ้น 344,373 ธุรกรรม โดยทั้งสองวิธีผู้วิจัยแบ่งสัดส่วนข้อมูลออกเป็น 80 ต่อ 20 โดยแบ่งเป็นชุดข้อมูลที่ใช้สอนโมเดล (Training Set) ร้อยละ 80 และชุดข้อมูลที่ใช้ทดสอบโมเดล (Test Set) ร้อยละ 20



ภาพที่ 3.5 สถาปัตยกรรมโครงสร้างของโมเดล Artificial Neural Network

จึงกล่าวได้ว่าระบบแนะนำสินค้าแบบผสม (A Hybrid Recommendation System) ที่ผู้วิจัยนำเสนอ นั้นสามารถใช้แก้ปัญหาการแนะนำสินค้าให้แก่ลูกค้ารายใหม่ได้ โดยหากมีลูกค้ารายใหม่เข้าสู่ระบบระบบจะแนะนำโดยอ้างอิงข้อมูลทางประชากรศาสตร์ร่วมกับการใช้ฟังก์ชันอายุ วิธีการเรียนรู้ของเครื่องรวมถึงการใช้บริบทร่วมด้วย แต่หากเป็นลูกค้ารายเก่าระบบจะสลับมาใช้วิธี Collaborative Filtering (CF) หรือ Artificial Neural Network (ANN) ตามความเหมาะสมของ

ข้อมูลและประสิทธิภาพของโมเดลที่ดีกว่า โดยขั้นตอนการแนะนำสินค้าให้แก่ลูกค้าทั้งกรณีลูกค้ารายเก่าและรายใหม่สามารถสรุปได้ดังภาพที่ 3.6



ภาพที่ 3.6 กรอบแนวคิดของระบบแนะนำสินค้าแบบผสม

บทที่ 4

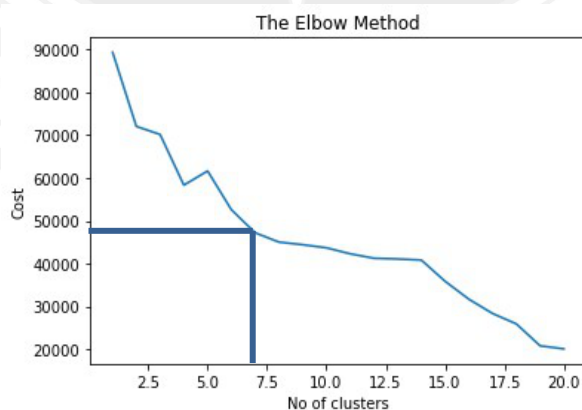
ผลการดำเนินการวิจัย

ในส่วนที่ 4 ผู้วิจัยนำเสนอผลการดำเนินการวิจัย โดยจะแบ่งออกเป็น 2 ส่วนหลักได้แก่ 1) ผลลัพธ์ของโมเดลแนะนำสินค้าสำหรับลูกค้ารายใหม่ และ 2) ผลลัพธ์ของโมเดลแนะนำสินค้าสำหรับลูกค้ารายเก่า

4.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่

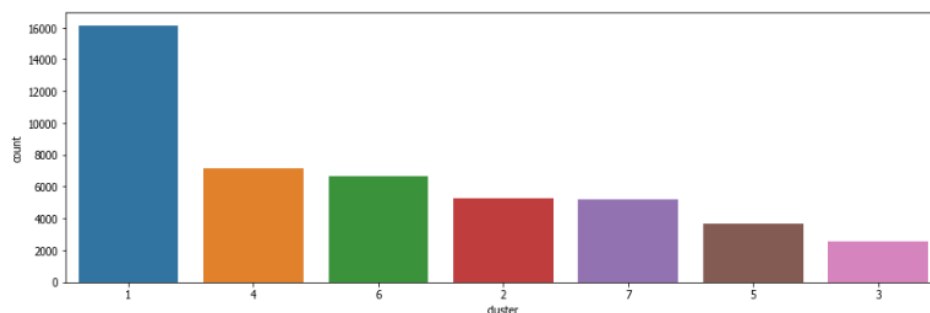
4.1.1 การแบ่งกลุ่มลูกค้ารายเก่า

จากการแบ่งกลุ่มลูกค้าเก่าในระบบตามลักษณะทางประชากรด้วยวิธี K-modes Clustering ผู้วิจัยกำหนดค่า K หรือจำนวนกลุ่มที่เหมาะสมในการจัดกลุ่ม โดยใช้วิธี Elbow Method ซึ่งจากกราฟจะพบว่าจุดที่กราฟเริ่มหักศอกหรือบริเวณที่ค่า Cost เริ่มลดลงอย่างคงที่หรือเพียงเล็กน้อยคือบริเวณที่ K=7 โดยเมื่อเพิ่มจำนวน K=8 ค่า Cost จะลดลงอีกเพียงเล็กน้อยเท่านั้น และจากกราฟจะพบว่าบริเวณที่เริ่มมีการหักศอกอีกครั้งคือเมื่อ K=19 แต่เนื่องจากตัวแปรที่ใช้จัดกลุ่มและจำนวนข้อมูลที่ไม่ได้มากนัก ผู้วิจัยจึงเลือกกำหนดค่า K ที่เหมาะสมเท่ากับ 7 กลุ่ม โดยแสดงกราฟ Elbow ไว้ในภาพที่ 4.1



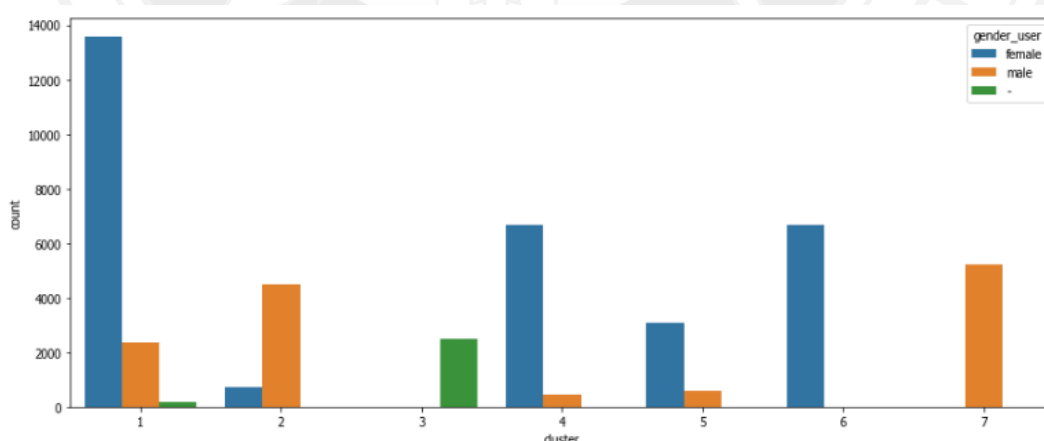
ภาพที่ 4.1 แสดงกราฟ Elbow Method

การแสดงผลข้อมูลเปรียบเทียบระหว่างกลุ่มของลูกค้าที่เป็นสมาชิกร้านอาหาร



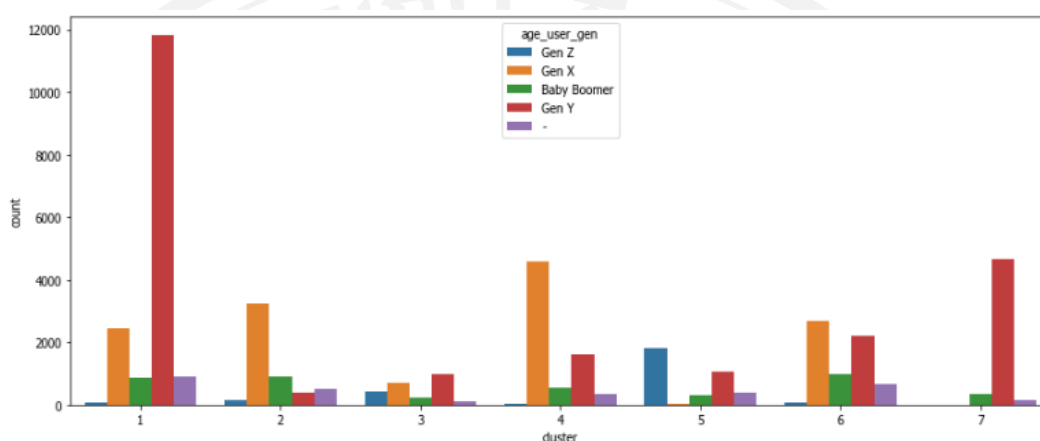
ภาพที่ 4.2 เปรียบเทียบจำนวนลูกค้าทั้ง 7 กลุ่ม

จากการแสดงผลข้อมูลในภาพที่ 4.2 พบว่าจากการแบ่งกลุ่มลูกค้าออกเป็น 7 กลุ่ม กลุ่มที่มีจำนวนสมาชิกมากที่สุดคือกลุ่มที่ 1 ถือเป็นฐานลูกค้าหลักของร้านอาหารโดยมีจำนวนคิดเป็นร้อยละ 34.8 จากจำนวนลูกค้าที่เป็นสมาชิกร้านอาหารทั้งหมด ส่วนลูกค้าในกลุ่มที่ 2-7 มีสัดส่วนจำนวนสมาชิกในแต่ละกลุ่มที่น้อยกว่ากลุ่มที่ 1 ถึง 2 เท่าตัว โดยแต่ละกลุ่มมีจำนวนสมาชิกโดยเฉลี่ยที่ใกล้เคียงกันซึ่งจะลดหลั่นกันไปตามภาพ โดยกลุ่มที่ควรทำการตลาดเพิ่มเพื่อขยายฐานลูกค้าให้มีจำนวนเพิ่มมากขึ้นในอนาคตคือลูกค้าในกลุ่มที่ 3 และ 5 เนื่องจากเป็นกลุ่มที่มีจำนวนสมาชิกน้อยกว่ากลุ่มอื่น ๆ โดยมีจำนวนลูกค้าคิดเป็นร้อยละ 6.5 และ 10.9 ตามลำดับ ซึ่งรายละเอียดลักษณะทางประชากรของลูกค้าในแต่ละกลุ่ม ผู้วิจัยจะอธิบายในลำดับต่อไป



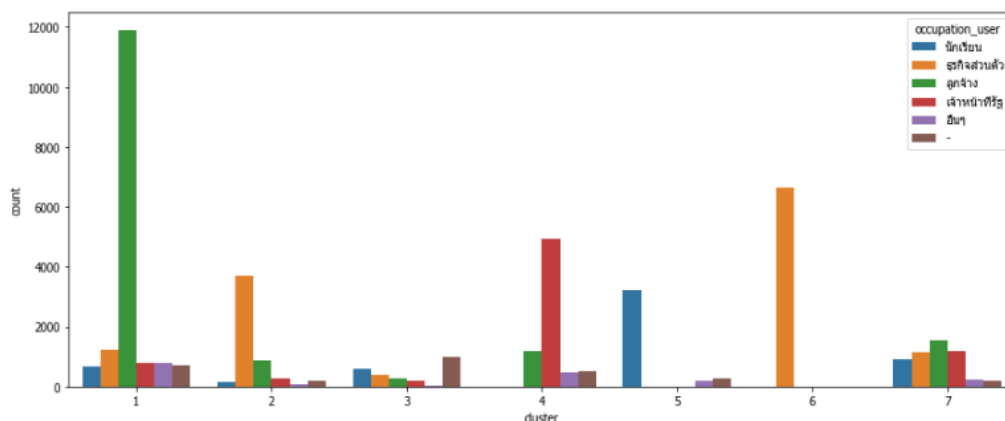
ภาพที่ 4.3 เปรียบเทียบจำนวนตัวแปรเพศของลูกค้าทั้ง 7 กลุ่ม

จากการแสดงข้อมูลในภาพที่ 4.3 ซึ่งเป็นการเปรียบเทียบตัวแปรเพศในแต่ละกลุ่มพบว่า ร้านลูกค้าหลักของร้านอาหารส่วนมากเป็นลูกค้าเพศหญิงเนื่องจากลูกค้าในกลุ่มที่ 1 ซึ่งเป็นกลุ่มที่มีจำนวนสมาชิกมากที่สุดมีลูกค้าที่เป็นเพศหญิงถึงร้อยละ 88.2 เช่นเดียวกับลูกค้าในกลุ่มที่ 4 และ 6 ซึ่งเป็นกลุ่มที่มีจำนวนสมาชิกรองลงมา ส่วนกลุ่มที่มีลูกค้าเพศชายเป็นหลักนั้นมีเพียง 2 กลุ่มจากทั้งหมด ได้แก่ลูกค้าในกลุ่มที่ 2 และ 7



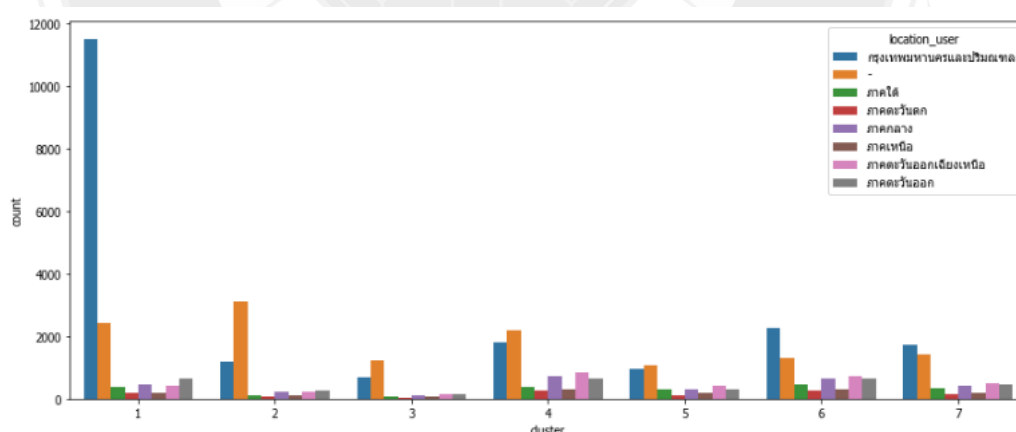
ภาพที่ 4.4 เปรียบเทียบจำนวนตัวแปรอายุของลูกค้าทั้ง 7 กลุ่ม

จากการแสดงข้อมูลในภาพที่ 4.4 ซึ่งเป็นการเปรียบเทียบตัวแปรอายุในแต่ละกลุ่มพบว่า ลูกค้าในกลุ่มที่ 1 ซึ่งเป็นกลุ่มลูกค้าหลักของร้านอาหารส่วนมากมีอายุอยู่ในกลุ่ม Gen Y หรือมีอายุในช่วง 22 - 39 ปี (เทียบอายุในปี พ.ศ. 2562) ส่วนลูกค้าในกลุ่มที่ 4 และ 6 ซึ่งเป็นกลุ่มที่มีจำนวนสมาชิกรองลงมาเป็นกลุ่มลูกค้าที่มีอายุในช่วง Gen X หรือมีอายุในช่วง 40 - 54 ปี (เทียบอายุในปี พ.ศ. 2562) ส่วนกลุ่มลูกค้าเพศชายหรือลูกค้าในกลุ่มที่ 2 และ 7 นั้นมีอายุอยู่ในช่วง Gen X และ Gen Y ตามลำดับ ส่วนกลุ่มที่ 5 เป็นกลุ่มที่ลูกค้าส่วนมากของกลุ่มอยู่ในช่วงอายุ Gen Z หรือมีอายุในช่วง 10 - 21 ปี (เทียบอายุในปี พ.ศ. 2562)



ภาพที่ 4.5 เปรียบเทียบจำนวนตัวแปรอาชีพของลูกค้าทั้ง 7 กลุ่ม

จากการแสดงข้อมูลในภาพที่ 4.5 ซึ่งเป็นการเปรียบเทียบตัวแปรอาชีพในแต่ละกลุ่มพบว่า ลูกค้าในกลุ่มที่ 1 ซึ่งเป็นกลุ่มลูกค้าหลักของร้านอาหารส่วนมากประกอบอาชีพลูกจ้าง ส่วนลูกค้าในกลุ่มที่ 4 และ 6 ซึ่งเป็นกลุ่มที่มีจำนวนสมาชิกรองลงมาส่วนมากประกอบอาชีพเจ้าหน้าที่รัฐ และธุรกิจส่วนตัวตามลำดับ เช่นเดียวกับลูกค้าในกลุ่มที่ 2 ที่ประกอบอาชีพธุรกิจส่วนตัวเป็นหลัก ส่วนลูกค้าในกลุ่มที่ 7 นั้นประกอบอาชีพหลากหลายกระจายกันไปเป็นจำนวนที่เท่า ๆ กันในแต่ละอาชีพ ลูกค้าในกลุ่มที่ 5 คือกลุ่มของนักเรียนซึ่งเป็นไปตามกลุ่มอายุที่ส่วนมากมีอายุอยู่ในช่วง Gen Z



ภาพที่ 4.6 เปรียบเทียบจำนวนตัวแปรภูมิลำเนาของลูกค้าทั้ง 7 กลุ่ม

จากการแสดงข้อมูลในภาพที่ 4.6 ซึ่งเป็นการเปรียบเทียบตัวแปรภูมิลำเนาในแต่ละกลุ่มพบว่าลูกค้าในกลุ่มที่ 1 ซึ่งเป็นกลุ่มลูกค้าหลักของร้านอาหารมากกว่าร้อยละ 80 อาศัยอยู่ในจังหวัดกรุงเทพมหานครและปริมณฑล ส่วนลูกค้าในกลุ่มอื่น ๆ นั้นมีการกระจายตัวกันไปในทุก ๆ ภูมิภาค

ดังนั้นจากภาพแสดงข้อมูลเปรียบเทียบตัวแปรประชากรที่แสดงในภาพที่ 4.2 – 4.6 สามารถสรุปได้ดังนี้ กลุ่มที่ 1 เป็นกลุ่มที่มีจำนวนสมาชิกมากที่สุดหรือฐานลูกค้าหลักของร้านอาหาร โดยลูกค้าส่วนใหญ่ในกลุ่มนี้เป็นลูกค้าเพศหญิงที่มีอายุอยู่ในช่วง Gen Y ประกอบอาชีพลูกจ้างและอาศัยอยู่ในเขตกรุงเทพมหานคร กลุ่มถัดมาซึ่งเป็นฐานลูกค้าที่รองลงมายังคงเป็นกลุ่มลูกค้าเพศหญิง ได้แก่ลูกค้าในกลุ่มที่ 4 และ 6 ตามลำดับ โดยมีอายุในช่วง Gen X ทั้งสองกลุ่ม โดยลูกค้าในกลุ่มที่ 4 จะเป็นลูกค้าที่ประกอบอาชีพเจ้าหน้าที่รัฐเป็นส่วนมาก ส่วนลูกค้าในกลุ่มที่ 6 เป็นกลุ่มที่ประกอบอาชีพธุรกิจส่วนตัว กลุ่มถัดมาคือกลุ่มของลูกค้าเพศชายซึ่งประกอบด้วยลูกค้าในกลุ่มที่ 2 และ 7 โดยมีอายุอยู่ในช่วง Gen X และ Gen Y ตามลำดับ โดยลูกค้าในกลุ่มที่ 2 จะประกอบอาชีพธุรกิจส่วนตัวเป็นหลัก ส่วนลูกค้าในกลุ่มที่ 7 นั้นประกอบอาชีพหลากหลายกระจายกันไปทั้งลูกจ้าง เจ้าหน้าที่รัฐ ธุรกิจส่วนตัวและนักเรียนในจำนวนที่เท่า ๆ กันในแต่ละอาชีพ ถัดมาคือลูกค้าที่มีอายุในช่วง Gen Z หรือกลุ่มที่เป็นนักเรียนนักศึกษาซึ่งมีจำนวนเพศหญิงต่อเพศชายในอัตราส่วนประมาณ 4 ต่อ 1 ซึ่งถือว่าเป็นกลุ่มที่ร้านอาหารมีจำนวนฐานลูกค้าน้อยและควรทำการตลาดเพื่อขยายฐานลูกค้าเพิ่มในอนาคต ส่วนในกลุ่มสุดท้ายซึ่งมีจำนวนสมาชิกน้อยที่สุดนั้นได้แก่ลูกค้ากลุ่มที่ 3 โดยเป็นกลุ่มลูกค้าที่ไม่ระบุเพศ

4.1.2 โมเดลต้นไม้ (Decision Tree)

ค่าความถูกต้องของโมเดล Decision Tree ที่ใช้ในการทำนายกลุ่มของลูกค้ารายใหม่เมื่อเข้าสู่ระบบมีค่าความถูกต้องที่ร้อยละ 100 ทั้งในกรณีของชุดข้อมูลสอน (Training Set) และชุดข้อมูลทดสอบ (Test Set) ซึ่งเป็นการพิสูจน์ว่าโมเดลที่ได้ไม่ได้เกิดปัญหา Over-fitting รวมทั้งผู้วิจัยได้ทำ Cross Validation เพื่อทดสอบโมเดลกับข้อมูลหลายชุดแต่ความถูกต้องของโมเดลนั้นก็ยังมีค่า 1.00 หรือร้อยละ 100 ในทุกกรณี โดยผลลัพธ์ของโมเดลสามารถแสดงได้ดังตารางที่ 4.2 ซึ่งเหตุผลที่ทำให้ค่าความถูกต้องมีความถูกต้องร้อยละ 100 นั้นผู้วิจัยมีข้อสันนิษฐานว่าเกิดจากการที่ตัวแปรที่ใช้ในการสร้างโมเดลต้นไม้มีจำนวนเพียง 4 ตัวแปรได้แก่ เพศ อายุ อาชีพ และภูมิลำเนา รวมถึงแต่ละตัวแปรนั้นเก็บค่าข้อมูลอยู่ในรูปแบบประเภทกลุ่มหรือ Category ทั้งหมด อีกทั้งตัวแปรแต่ละตัวมีจำนวน Class ไม่มากนัก จึงอาจส่งผลให้ขั้นตอนในการสร้างโมเดลนั้นสามารถสร้างกฎในการจำแนกกลุ่มได้ครอบคลุมทั้งหมดทุกกรณี

โดยค่าที่ใช้วัดความถูกต้องของโมเดลต้นไม้มีสูตรในการคำนวณดังนี้

$$\text{Precision} = \frac{TP}{TP + FP}$$

$$\text{Recall} = \frac{TP}{TP + FN}$$

$$F1\text{-score} = \frac{2TP}{2TP + FP + FN}$$

เมื่อ TP = True Positive, TN = True Negative

FP = False Positive, FN = False Negative

โดยเป็นค่าที่ได้จากตาราง Confusion Matrix ดังแสดงในตารางที่ 4.1

ตารางที่ 4.1 ตาราง Confusion Matrix ของโมเดลต้นไม้

กลุ่มของลูกคำที่โมเดลแนะนำ			
กลุ่มที่แท้จริง ของลูกคำ	Positive		Negative
	Positive	TP	FN
	Negative	FP	TN

ตารางที่ 4.2 ผลลัพธ์ค่าความถูกต้องของโมเดลต้นไม้

Cluster	Precision	Recall	f1-score
1	1.00	1.00	1.00
2	1.00	1.00	1.00
3	1.00	1.00	1.00
4	1.00	1.00	1.00
5	1.00	1.00	1.00
6	1.00	1.00	1.00
7	1.00	1.00	1.00

4.1.3 โมเดลแนะนำสินค้าสำหรับลูกค้ารายใหม่

โมเดลแนะนำสินค้าสำหรับลูกค้ารายใหม่นั้นประกอบด้วยการนำ Demographic-based หรือข้อมูลทางประชากรศาสตร์มาใช้ร่วมกับฟังก์ชันอายุ (Similarity Age Function) และบริบท (Contexts) โดยโมเดลแนะนำรายการอาหารสำหรับลูกค้ารายใหม่แบ่งออกเป็นสองกรณีดังนี้

- 1) โมเดลแนะนำรายการอาหารทั้ง 80 เมนู (ครบทุกรายการอาหารของร้านอาหาร)
- 2) โมเดลแนะนำรายการอาหารเฉพาะเมนูขายดี 10 เมนูแรก

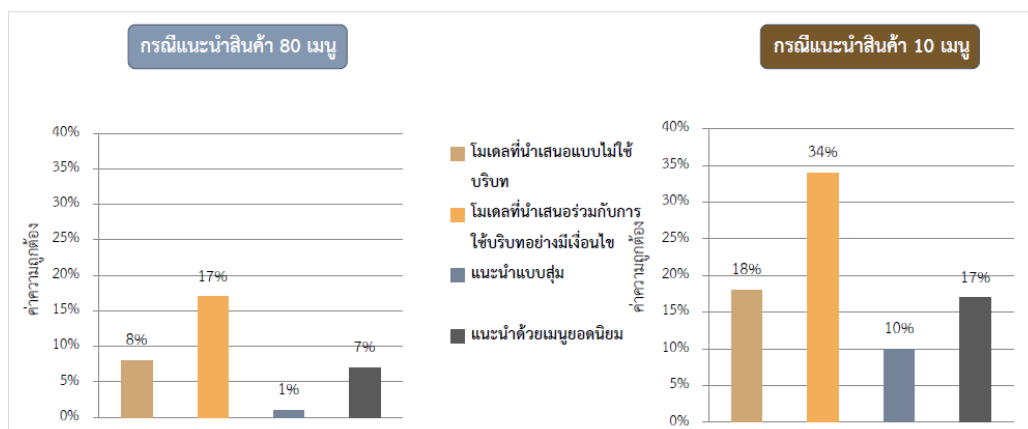
โดยในแต่ละกรณีจะเปรียบเทียบความถูกต้องในการแนะนำของ 4 วิธีดังนี้ 1) โมเดลที่นำเสนอร่วมกับการใช้บริบทอย่างมีเงื่อนไข 2) โมเดลที่นำเสนอแบบไม่ใช้บริบท 3) การแนะนำแบบสุ่ม 4) การแนะนำด้วยเมนูยอดนิยม

ผลการดำเนินงานพบว่ากรณีของโมเดลแนะนำรายการอาหารทั้ง 80 เมนู โมเดลที่ผู้วิจัยนำเสนอร่วมกับการใช้บริบทอย่างมีเงื่อนไขมีประสิทธิภาพในการแนะนำรายการอาหารให้แก่ลูกค้ารายใหม่ได้ดีที่สุดโดยมีค่าความถูกต้องที่ร้อยละ 17 สูงกว่าวิธีการแนะนำแบบสุ่มและการแนะนำด้วยเมนูยอดนิยมที่ร้อยละ 16 และ 10 ตามลำดับ รวมถึงดีกว่ากรณีที่ไม่มีการใช้บริบทร่วมด้วยถึงร้อยละ 9 ส่วนในกรณีของโมเดลแนะนำรายการอาหารเฉพาะเมนูขายดี 10 เมนูแรก ผลปรากฏว่าโมเดลที่ผู้วิจัยนำเสนอร่วมกับการใช้บริบทอย่างมีเงื่อนไขมีประสิทธิภาพในการแนะนำรายการอาหารดีขึ้นหนึ่งเท่าตัวโดยมีค่าความถูกต้องเพิ่มขึ้นเป็นร้อยละ 34 ส่วนโมเดลอื่น ๆ นั้นมีค่าความถูกต้องเพิ่มขึ้นที่ประมาณร้อยละ 10

ดังนั้นกล่าวได้ว่า โมเดลที่ผู้วิจัยนำเสนอกรณีที่ใช้ร่วมกับบริบทมีประสิทธิภาพในการแนะนำรายการอาหารให้แก่ผู้ใช้รายใหม่ได้ดีที่สุดจากทั้ง 4 วิธี ซึ่งมีความสอดคล้องกับข้อสมมติฐานเบื้องต้นที่กล่าวว่า บริบทนั้นมีผลต่อการตัดสินใจซื้อสินค้าของผู้ใช้และการนำบริบทมาใช้ร่วมกับโมเดลแนะนำสินค้านั้นจะช่วยเพิ่มประสิทธิภาพของระบบแนะนำสินค้าได้ โดยผลลัพธ์ทั้ง 2 กรณีที่ผู้วิจัยนำเสนอไว้ในตารางที่ 4.3 และแสดงกราฟเปรียบเทียบประสิทธิภาพของทั้ง 4 วิธีดังภาพที่ 4.7

ตารางที่ 4.3 ตารางเปรียบเทียบประสิทธิภาพการทำนายทั้ง 4 วิธี

วิธี	ค่าความถูกต้อง	
	รายการอาหาร 80 เมนู	รายการอาหาร 10 เมนู
1 โมเดลที่นำเสนอแบบไม่ใช้บริบท	ร้อยละ 8	ร้อยละ 18
2 โมเดลที่นำเสนอร่วมกับการใช้บริบทอย่างมีเงื่อนไข	ร้อยละ 17	ร้อยละ 34
3 แนะนำแบบสุ่ม	ร้อยละ 1	ร้อยละ 10
4 แนะนำด้วยเมนูยอดนิยม	ร้อยละ 7	ร้อยละ 17



ภาพที่ 4.7 กราฟแสดงการเปรียบเทียบประสิทธิภาพของโมเดลทั้ง 2 กรณี

โดยค่าความถูกต้องมีสูตรในการคำนวณดังนี้

$$\text{ความถูกต้อง} = \frac{TP + TN}{TP + TN + FN + FP}$$

เมื่อ TP = True Positive, TN = True Negative

FP = False Positive, FN = False Negative

โดยเป็นค่าที่ได้จากตาราง Confusion Matrix ดังแสดงในตารางที่ 4.4

ตารางที่ 4.4 ตาราง Confusion Matrix ของโมเดลแนะนำรายการอาหาร

รายการอาหารที่โมเดลแนะนำ			
รายการอาหาร ที่ลูกค้าสั่งจริง			
	Positive	Negative	
	Positive	TP	FN
	Negative	FP	TN

4.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า

ผลลัพธ์ของระบบแนะนำสินค้าสำหรับลูกค้ารายเก่าแบ่งออกเป็น 2 กรณีได้แก่ กรณีที่ไม่มีการใช้บริบทและกรณีที่มีการใช้บริบทอย่างมีเงื่อนไขร่วมกับโมเดล โดยในแต่ละกรณีจะแบ่งออกเป็น

2 วิธีได้แก่ วิธี Collaborative Filtering และ วิธี Neural Network และทั้งสองวิธีนี้จะใช้กับการแนะนำเมนูอาหารครบทั้ง 80 รายการของร้านอาหาร

จากผลการดำเนินงานพบว่ากรณีที่ไม่มีการใช้บริบท โมเดล Neural Network มีประสิทธิภาพในการทำนาย Rating ต่ำกว่าวิธี Collaborative Filtering โดยมีค่าความผิดพลาดในการทำนายคะแนนความชื่นชอบหรือค่า RMSE ของโมเดล Neural Network ต่ำกว่าค่า RMSE ของโมเดล Collaborative Filtering เท่ากับ 0.14

ส่วนในกรณีที่มีการใช้บริบทร่วมกับโมเดลนั้น โมเดลที่มีประสิทธิภาพในการทำนาย Rating ต่ำกว่าคือโมเดล Neural Network โดยมีค่าความผิดพลาดในการทำนายคะแนนความชื่นชอบหรือค่า Average RMSE เท่ากับ 1.22 และ 2.02 สำหรับชุดข้อมูลเรียนรู้และชุดข้อมูลทดสอบตามลำดับ ซึ่งเป็นโมเดลที่มีค่าความผิดพลาดที่ต่ำที่สุดจากทั้ง 4 โมเดล และโมเดล Collaborative Filtering กรณีที่มีการใช้บริบทร่วมกับโมเดลมีค่า Average RMSE สูงกว่ากรณีที่ไม่มีการใช้บริบทซึ่งไม่เป็นไปตามข้อสมมติฐานของผู้วิจัยที่ว่า การนำบริบทมาใช้ร่วมกับโมเดลจะช่วยเพิ่มประสิทธิภาพการทำนายคะแนนความชื่นชอบให้ดียิ่งขึ้น โดยผู้วิจัยมีข้อสันนิษฐานว่าอาจเกิดจากการแบ่งข้อมูลออกตามบริบทที่กำหนดแล้วส่งผลให้ข้อมูลที่ค่อนข้างว่างอยู่แล้วเกิดมีจำนวนว่างเพิ่มมากขึ้น จึงส่งผลให้โมเดล Collaborative Filtering กรณีที่มีการใช้บริบทมีประสิทธิภาพในการทำนาย Rating ที่แย่ลง

ดังนั้นจึงกล่าวได้ว่าโมเดล Neural Network มีประสิทธิภาพในการทำนาย Rating ต่ำกว่าโมเดล Collaborative Filtering ทั้งในกรณีที่ไม่มีการใช้บริบทและใช้บริบทร่วมกับโมเดล นอกจากนั้นการนำบริบทมาใช้ร่วมกับโมเดลยังเป็นการช่วยเพิ่มประสิทธิภาพในการทำนาย Rating ของลูกค้ารายเก่าให้ดียิ่งขึ้นได้ซึ่งสอดคล้องกับงานวิจัยที่ผู้วิจัยทบทวนวรรณกรรมไว้ในบทที่ 2

โดยค่าประสิทธิภาพการทำนาย Rating ของทุกกรณีผู้วิจัยได้นำเสนอไว้ในตารางที่ 4.5 ซึ่งประสิทธิภาพการทำนาย Rating จะพิจารณาจากค่า RMSE และ Average RMSE ดังนี้

ค่า RMSE (Root Mean Square Error) หรือค่าความผิดพลาดในการทำนายคะแนนความชื่นชอบของลูกค้าโดยเปรียบเทียบระหว่างคะแนนความชื่นชอบที่โมเดลทำนายกับคะแนนความชื่นชอบที่แท้จริงของลูกค้าในแต่ละเมนู โดยมีสูตรในการคำนวณดังนี้

$$RMSE = \sqrt{\frac{1}{n} \sum_{u,i} (\hat{r}_{u,i} - r_{u,i})^2}$$

เมื่อ $\hat{r}_{u,i}$ คือ ค่า Rating ที่โมเดลทำนายของลูกค้า u ที่เมนู i

$r_{u,i}$ คือ ค่า Rating ที่แท้จริงของลูกค้า u ที่เมนู i

n คือ จำนวนการทำนาย Rating ทั้งหมดในชุดข้อมูล

ค่า Average RMSE (Average Root Mean Square Error) หรือค่าความผิดพลาดในการทำนายคะแนนความชื่นชอบของลูกค้าโดยเปรียบเทียบระหว่างคะแนนความชื่นชอบที่โมเดลทำนายกับคะแนนความชื่นชอบที่แท้จริงของลูกค้าในแต่ละเมนู และนำมาเฉลี่ยจากจำนวน Cross Contexts ทั้งหมด โดยมีสูตรในการคำนวณดังนี้

$$\text{Average RMSE} = \frac{\sum_C \sqrt{\frac{1}{n} \sum_{u,i} (\hat{r}_{u,i} - r_{u,i})^2}}{C}$$

เมื่อ $\hat{r}_{u,i}$ คือ ค่า rating ที่โมเดลทำนายของลูกค้า u ที่เมนู i

$r_{u,i}$ คือ ค่า rating ที่แท้จริงของลูกค้า u ที่เมนู i

n คือ จำนวนการทำนาย rating ทั้งหมดในชุดข้อมูล

C คือ จำนวน cross contexts ทั้งหมดที่เป็นไปได้

ตารางที่ 4.5 ประสิทธิภาพการทำนาย Rating ของแต่ละโมเดล

กรณี	โมเดล	RMSE		Average RMSE	
		train	test	train	test
ไม่ใช้บริบท	Collaborative Filtering	1.49	2.18	-	-
	Neural Network	1.34	2.04	-	-
ใช้บริบท	Collaborative Filtering	-	-	1.51	2.30
	Neural Network	-	-	1.22	2.02

บทที่ 5

สรุป อภิปรายผล และข้อเสนอแนะ

ระบบแนะนำสินค้า (Recommendation System: RS) เป็นระบบที่ช่วยกระตุ้นยอดขายและมอบประสบการณ์ที่ดีในการได้รับบริการของลูกค้า โดยจะแนะนำสินค้าที่คาดว่าผู้ใช้น่าจะชื่นชอบโดยอ้างอิงจากข้อมูลการสั่งซื้อสินค้าในอดีตของลูกค้าหรือคะแนนความชื่นชอบที่ได้รับจากลูกค้าโดยตรง (Explicit Rating) เพื่อเพิ่มโอกาสที่ผู้ใช้จะซื้อสินค้ามากขึ้นทำให้ระบบแนะนำมีบทบาทอย่างมากในแวดวงธุรกิจ แต่หนึ่งในปัญหาที่สำคัญที่พบในระบบแนะนำสินค้าแบบดั้งเดิมคือปัญหาผู้ใช้งานระบบรายใหม่ (Cold-start Problem) โดยเป็นปัญหาในกรณีที่ระบบไม่เคยมีข้อมูลคะแนนความชื่นชอบในอดีตของผู้ใช้รายใหม่มาก่อน จึงทำให้ไม่สามารถแนะนำสินค้าที่คาดว่าผู้ใช้รายใหม่จะชื่นชอบได้

ผู้วิจัยนำเสนอวิธีการแนะนำสินค้าแบบผสมเพื่อใช้ในการแก้ปัญหาผู้ใช้งานระบบรายใหม่และฟังก์ชันการคิดคะแนนความชื่นชอบโดยนัยของลูกค้ารายเก่า (Implicit Rating Function) ซึ่งนำเสนอวิธีที่เหมาะสมกับธุรกิจร้านอาหารและสามารถนำไปประยุกต์ใช้กับข้อมูลที่มีอยู่ได้จริง โดยแบ่งการสรุปผลออกเป็นสามส่วนดังนี้

- 5.1 สรุปและอภิปรายผล
- 5.2 ข้อจำกัดและปัญหาที่พบในงานวิจัย
- 5.3 ข้อเสนอแนะในอนาคต

5.1 สรุปและอภิปรายผล

5.1.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่

ผลการดำเนินงานพบว่าประสิทธิภาพของโมเดลที่นำเสนอสำหรับลูกค้ารายใหม่ร่วมกับการใช้บริบทอย่างมีเงื่อนไขในกรณีโมเดลแนะนำรายการอาหารทั้ง 80 เมนูมีความถูกต้องในการแนะนำรายการอาหารสูงกว่าวิธีการแนะนำแบบสุ่มและการแนะนำด้วยเมนูยอดนิยมที่ร้อยละ 16 และ 10 ตามลำดับ รวมถึงดีกว่ากรณีที่ไม่มีการใช้บริบทร่วมด้วยถึงร้อยละ 9 แต่เนื่องจากข้อมูลที่ใช้ในการวิเคราะห์มีจำนวนว่างค่อนข้างมาก (Sparse Data) ส่งผลให้ค่าความถูกต้องในการทำนายของ

โมเดลไม่สูงมากนัก โดยเมื่อทำการเปรียบเทียบค่าความถูกต้องของโมเดลแนะนำรายการอาหารกรณีแนะนำรายการอาหารเฉพาะเมนูขายดี 10 เมนูแรกของร้าน พบว่าโมเดลที่ผู้วิจัยนำเสนอร่วมกับการใช้บริบทนั้นมีค่าความถูกต้องในการทำนายเพิ่มสูงขึ้นถึงหนึ่งเท่าตัวโดยมีค่าความถูกต้องเพิ่มขึ้นเป็นร้อยละ 34 และยังคงดีกว่าวิธีการแนะนำแบบอื่น ๆ ทั้ง 3 วิธี โดยมีค่าความถูกต้องในการแนะนำรายการอาหารสูงกว่าวิธีการแนะนำแบบสุ่มและการแนะนำด้วยเมนูยอดนิยมที่ร้อยละ 24 และ 17 ตามลำดับ รวมถึงดีกว่ากรณีที่ไม่มีการใช้บริบทร่วมด้วยถึงร้อยละ 16

ดังนั้นโมเดลที่ผู้วิจัยนำเสนอสามารถแนะนำรายการอาหารให้แก่ผู้ใช้รายใหม่ได้ดีกว่าวิธีการแนะนำแบบดั้งเดิมทั้งสองวิธีคือวิธีการแนะนำแบบสุ่มและวิธีการแนะนำด้วยเมนูยอดนิยมในทั้งสองกรณี โดยโมเดลที่ผู้วิจัยนำเสนอร่วมกับการใช้บริบทมีค่าความถูกต้องในการทำนายดีกว่ากรณีที่ไม่มีการใช้บริบทร่วมด้วย และเมื่อนำโมเดลมาใช้แนะนำรายการอาหารในจำนวนที่น้อยลงโมเดลจะมีประสิทธิภาพในการทำนายดีขึ้นในทุก ๆ กรณี ซึ่งโดยทั่วไปแล้วระบบแนะนำสินค้าส่วนใหญ่มักจะใช้ในการแนะนำสินค้าเพียงแค่ 10 หรือ 5 เมนูขายดีเท่านั้น จึงกล่าวได้ว่าโมเดลที่ผู้วิจัยเสนอนั้นสามารถแก้ปัญหาผู้ใช้งานระบบรายใหม่ได้

5.1.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า

ผู้วิจัยนำเสนอฟังก์ชันการคิดคะแนนความชื่นชอบโดยนัย (Implicit Rating Function) ซึ่งเป็นวิธีที่เหมาะสมกับธุรกิจร้านอาหารและสามารถนำไปประยุกต์ใช้กับข้อมูลที่มีอยู่ได้จริง จากนั้นจะนำค่าที่ได้ไปใช้ในการสร้างโมเดลในลำดับต่อไป สำหรับโมเดลในการแนะนำอาหารสำหรับลูกค้ารายเก่าผู้วิจัยได้นำเสนอ 2 วิธีหลัก ได้แก่วิธี Collaborative Filtering และ Neural Network แม้ว่าโมเดล Neural Network with Contexts จะมีประสิทธิภาพในการทำนาย Rating ดีที่สุดจากทั้ง 4 โมเดล แต่เนื่องจากค่าความผิดพลาดในการทำนายที่แตกต่างกันเพียง 0.28 ดังนั้นสำหรับในปัจจุบันโมเดล Neural Network อาจจะยังไม่คุ้มต่อการนำไปใช้เมื่อเทียบกับวิธี Collaborative Filtering ซึ่งเป็นโมเดลที่ทำได้ง่ายกว่ารวมถึงสามารถอธิบายหลักการทำงานได้ง่ายกว่าเมื่อเทียบกับวิธี Neural Network ที่เป็น Black Box ดังนั้นผู้วิจัยจึงนำเสนอโมเดลแนะนำสินค้าสำหรับลูกค้ารายเก่าไว้ทั้ง 2 วิธีเพื่อเป็นแนวทางให้องค์กรสามารถนำไปพัฒนาต่อยอดได้ในอนาคต

ดังนั้นจากทั้ง 2 กรณีสรุปได้ว่า โมเดลที่ดีที่สุดในการแนะนำรายการอาหารให้แก่ลูกค้ารายใหม่และลูกค้ารายเก่า คือโมเดลที่มีการนำบริบทมาใช้ร่วมด้วย จึงกล่าวได้ว่าการนำบริบทมาใช้ร่วมกับโมเดลนั้นถือเป็นการช่วยเพิ่มประสิทธิภาพให้แก่ระบบแนะนำสินค้าได้จริงดังที่ผู้วิจัยได้ตั้งข้อสมมติฐานไว้

5.2 ข้อจำกัดและปัญหาที่พบในงานวิจัย

5.2.1 ข้อจำกัดด้านข้อมูล

เนื่องจากชุดข้อมูลที่นำมาวิเคราะห์ในงานวิจัยเป็นชุดข้อมูลจริง ดังนั้นข้อจำกัดของงานวิจัยคือเรื่องของการเข้าถึงข้อมูลที่ทำได้อย่างจำกัด และข้อมูลที่จะนำมาวิเคราะห์นั้นมีเฉพาะในส่วนที่ร้านอาหารทำการเก็บข้อมูลไว้อยู่แล้วเท่านั้นไม่สามารถเก็บข้อมูลอื่น ๆ เพิ่มเติมได้ ยกตัวอย่างเช่นข้อมูลทางประชากรหรือข้อมูลส่วนตัวของลูกค้าที่เป็นสมาชิกร้านอาหาร โดยข้อมูลในส่วนนี้ร้านอาหารจะได้รับจากลูกค้าในขั้นตอนของการสมัครบัตรสมาชิกในขั้นตอนแรกเท่านั้นและไม่สามารถขอข้อมูลเพิ่มเติมอื่น ๆ ที่หลังได้ รวมถึงข้อมูลที่ลูกค้ากรอกให้แก่ร้านอาหารนั้นมี Missing Data ค่อนข้างมาก หรือลูกค้าบางท่านก็ไม่ได้ให้ข้อมูลส่วนตัว ซึ่งจะส่งผลให้ผู้วิจัยมีจำนวนข้อมูลที่จะนำมาใช้ในการวิเคราะห์น้อยลง อีกหนึ่งปัญหาที่สำคัญคือเรื่องของชุดข้อมูลการสั่งซื้อสินค้าของลูกค้าที่นำมาวิเคราะห์มีจำนวนว่างค่อนข้างมาก (Sparse Data) เนื่องจากการสั่งอาหารของลูกค้าไม่มีการกระจายตัวอย่างทั่วถึงในทุกเมนูอาหารเท่าที่ควร ซึ่งเป็นส่วนสำคัญที่ส่งผลต่อค่าความถูกต้องของระบบแนะนำสินค้าทั้งในส่วนของการทำนายเมนูให้แก่ลูกค้ารายใหม่และโมเดลทำนายคะแนนความชื่นชอบสำหรับลูกค้ารายเก่า

5.2.2 ข้อจำกัดของระบบแนะนำสินค้า

ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่ โมเดลที่ผู้วิจัยนำเสนอไม่เหมาะกับการนำไปใช้กับข้อมูลแบบ Real Time เนื่องจากโมเดลจะต้องทำการทำนายกลุ่มของลูกค้ารายใหม่เสียก่อนจึงจะสามารถแนะนำรายการอาหารที่คาดว่าลูกค้ารายใหม่จะชื่นชอบได้ ผู้วิจัยจึงมีความเห็นว่าร้านอาหารอาจจะใช้วิธีการออกแคมเปญส่งคูปอนให้แก่ลูกค้าใหม่ที่เพิ่งเป็นสมาชิกของร้าน โดยนำโมเดลมาทำนายว่าลูกค้าน่าจะชอบเมนูใด จากนั้นจึงส่งคูปอนไปให้ทางอีเมลเพื่อเป็นการเชิญชวนให้ลูกค้ามารับประทานอาหารที่ร้านมากขึ้น

5.3 ข้อเสนอแนะในอนาคต

5.3.1 ระบบแนะนำสินค้าสำหรับลูกค้ารายใหม่

จากผลลัพธ์ของโมเดลที่ใช้ในการแนะนำเมนูอาหารสำหรับลูกค้ารายใหม่นั้นพบว่าโมเดลอาจจะยังมีค่าความถูกต้องที่ไม่สูงมากนัก ดังนั้นสำหรับในอนาคตหากต้องการพัฒนาต่อยอดเพื่อเพิ่มประสิทธิภาพการทำนายให้มีค่าความถูกต้องของโมเดลเพิ่มสูงขึ้น อาจทำการเก็บข้อมูลด้านประชากรจากลูกค้าเพิ่มขึ้นเพื่อเพิ่มตัวแปรในการแบ่งกลุ่มลูกค้าให้ดียิ่งขึ้น เช่น ตัวแปรรายได้ ตัวแปร

สถานภาพ หรืออาจจะเพิ่มการเก็บข้อมูลของลูกค้าจากช่องทางอื่น เช่น เก็บข้อมูลพฤติกรรมการใช้งานของลูกค้าบนเว็บไซต์ร้านค้า เช่น การเลือกดูรายการอาหารต่าง ๆ ของลูกค้าใหม่บนเว็บไซต์ของทางร้าน ซึ่งในปัจจุบันยังไม่มีเก็บข้อมูลประเภทนี้ โดยสามารถนำข้อมูลมาวิเคราะห์เพิ่มเติมและช่วยในการแนะนำสินค้าให้ดียิ่งขึ้น

5.3.2 ระบบแนะนำสินค้าสำหรับลูกค้ารายเก่า

แม้โมเดลที่มีประสิทธิภาพดีที่สุดคือโมเดล Neural Network with Contexts แต่เนื่องจากประสิทธิภาพการทำนายระหว่างสองวิธีมีค่าความผิดพลาดที่แตกต่างกันไม่มากนัก ผู้วิจัยจึงแนะนำให้เลือกใช้วิธี Collaborative Filtering แทน แต่สำหรับในอนาคตหากองค์กรมีข้อมูลในปริมาณเพิ่มมากขึ้นและต้องการพัฒนาความถูกต้องของโมเดล วิธี Neural Network ถือเป็นทางเลือกที่ดี เนื่องจากเป็นวิธีที่ยังสามารถพัฒนาต่อไปได้ต่างจากวิธี Collaborative Filtering ที่อาจจะพัฒนาต่อไปได้อีกไม่มากนัก

บรรณานุกรม

- Anand, S. S., & Mobasher, B. (2007). *Contextual recommendation*. Paper presented at the From Web to Social Web: Discovering and Deploying User and Content Profiles, Berlin, Heidelberg.
- Bhatia, A. (2016, 8-10 Dec.). *Community detection for cold start problem in personalization: Community detection in large social network graphs based on users' structural similarities and their attribute similarities*. Paper presented at the 2016 IEEE International Conference on Computer and Information Technology (CIT).
- Bisogni, C. A., Falk, L. W., Madore, E., Blake, C. E., Jastran, M., Sobal, J., & Devine, C. M. (2007). Dimensions of everyday eating and drinking episodes. *Appetite, 48*(2), 218-231. doi:10.1016/j.appet.2006.09.004
- Bobadilla, J., Alonso, S., & Hernando, A. (2020). Deep learning architecture for collaborative filtering recommender systems. *Applied Sciences, 10*, 2441. doi:10.3390/app10072441
- Bobadilla, J., Ortega, F., Hernando, A., & Gutiérrez, A. (2013). Recommender systems survey. *Knowledge-Based Systems, 46*(C), 109-132. doi:10.1016/j.knosys.2013.03.012
- Braunhofer, M., Elahi, M., & Ricci, F. (2015). *User personality and the new user problem in a context-aware point of interest recommender system*. Paper presented at the Information and Communication Technologies in Tourism 2015, Cham.
- Chu, W.-T., & Tsai, Y.-L. (2017). A hybrid recommendation system considering visual information for predicting favorite restaurants. *World Wide Web, 20*. doi:10.1007/s11280-017-0437-1
- Covington, P., Adams, J., & Sargin, E. (2016). *Deep neural networks for YouTube recommendations*. Paper presented at the Proceedings of the 10th ACM Conference on Recommender Systems, Boston, Massachusetts, USA. Retrieved from <https://doi.org/10.1145/2959100.2959190>
- Fernández-Tobías, I., Braunhofer, M., Elahi, M., Ricci, F., & Cantador, I. (2016). Alleviating

the new user problem in collaborative filtering by exploiting personality information. *User Modeling and User-Adapted Interaction*, 26(2-3), 221-255. doi:10.1007/s11257-016-9172-z

Haoxian, F. (2017). *User behavior learning in designing restaurant recommender systems: An approach to leveraging historical data and implicit feedback*. (Master of Applied Science). University of Ottawa,

He, X., Liao, L., Zhang, H., Nie, L., Hu, X., & Chua, T.-S. (2017). *Neural collaborative filtering*. Paper presented at the Proceedings of the 26th International Conference on World Wide Web, Perth, Australia. Retrieved from <https://doi.org/10.1145/3038912.3052569>

Hu, R., & Pu, P. (2011). *Enhancing collaborative filtering systems with personality information*. Paper presented at the Proceedings of the fifth ACM conference on Recommender systems, Chicago, Illinois, USA. Retrieved from <https://doi.org/10.1145/2043932.2043969>

Kim, H.-r., & Chan, P. K. (2005, 26-28 May.). *Implicit indicators for interesting web pages*. Paper presented at the the Proceedings of the 1st International Conference on Web Information Systems and Technologies (WEBIST). Miami, USA.

Klein, N. M., & Yadav, M. S. (1989). Context effects on effort and accuracy in choice: An enquiry into adaptive decision making. *Journal of Consumer Research*, 15(4), 411-421. doi:10.1086/209181

Lee, T. Q., Park, Y., & Park, Y.-T. (2008). A time-based approach to effective recommender systems using implicit feedback. *Expert Systems With Applications*, 34(4), 3055-3062. doi:10.1016/j.eswa.2007.06.031

Lieberman, H., & Selker, T. (2000). Out of context: Computer systems that adapt to, and learn from, context. *IBM Systems Journal*, 39(3.4), 617-632. doi:10.1147/sj.393.0617

Lilien, G. L. (1992). *Marketing models*. Englewood Cliffs, N.J.: Prentice Hall.

Lussier, D. A., & Olshavsky, R. W. (1979). Task complexity and contingent processing in brand choice. *Journal of Consumer Research*, 6(2), 154-165. doi:10.1086/208758

Nair, P., Moh, M., & Moh, T. (2016, 31 Oct.-4 Nov.). *Using social media presence for alleviating cold start problems in privacy protection*. Paper presented at the

- 2016 International Conference on Collaboration Technologies and Systems (CTS).
- Oku, K., Nakajima, S., Miyazaki, J., & Uemura, S. (2006, 10-12 May). *Context-aware SVM for context-dependent information recommendation*. Paper presented at the 7th International Conference on Mobile Data Management (MDM'06).
- Oord, A. v. d., Dieleman, S., & Schrauwen, B. (2013). *Deep content-based music recommendation*. Paper presented at the Proceedings of the 26th International Conference on Neural Information Processing Systems - Volume 2, Lake Tahoe, Nevada.
- Pandey, A. K., & Rajpoot, D. S. (2016, 26-28 Dec.). *Resolving cold start problem in recommendation system using demographic approach*. Paper presented at the 2016 International Conference on Signal Processing and Communication (ICSC).
- Panniello, U., Tuzhilin, A., Gorgoglione, M., Palmisano, C., & Pedone, A. (2009). *Experimental comparison of pre- vs. post-filtering approaches in context-aware recommender systems*. Paper presented at the Proceedings of the third ACM conference on Recommender systems, New York, USA. Retrieved from <https://doi.org/10.1145/1639714.1639764>
- Parhi, P., Pal, A., & Aggarwal, M. (2017, 19-20 Jan.). *A survey of methods of collaborative filtering techniques*. Paper presented at the 2017 International Conference on Inventive Systems and Control (ICISC).
- Pazzani, M. (1999). A framework for collaborative, content-based and demographic filtering. *Artificial Intelligence Review*, 13(5), 393-408.
doi:10.1023/A:1006544522159
- Rana, C. A. Q., Salima, H., Usama, F., & Hammam, C. (2014, 23-25 June). *From a "cold" to a "warm" start in recommender systems*. Paper presented at the 2014 IEEE 23rd International WETICE Conference.
- Rentfrow, P. J., & Gosling, S. D. (2003). The Do Re Mi's of everyday life: The structure and personality correlates of music preferences. *Journal of Personality and Social Psychology*, 84(6), 1236-1256. doi:10.1037/0022-3514.84.6.1236
- Robertson, T. S., & Kassarian, H. H. (1991). *Handbook of consumer behavior*. Englewood Cliffs, N.J.: Prentice-Hall.

Schedl, M., Zamani, H., Chen, C.-W., Deldjoo, Y., & Elahi, M. (2018). Current challenges and visions in music recommender systems research. *International Journal of Multimedia Information Retrieval*, 7(2), 95-116. doi:10.1007/s13735-018-0154-2



ประวัติผู้เขียน

ชื่อ-นามสกุล

นางสาวพิมพ์ชนก ปทุมชาติ

ประวัติการศึกษา

วิทยาศาสตร์บัณฑิต

มหาวิทยาลัยธรรมศาสตร์

ปีสำเร็จการศึกษา พ.ศ. 2560

