



การเพิ่มศักยภาพระบบแนะนำด้วยการผสานการวิเคราะห์ความรู้สึกด้วยแบบจำลอง SieBERT เข้า
กับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้อย่างมีประสิทธิภาพ

SENTIMENT ENHANCED CF: UTILIZING SIEBERT FOR OPTIMIZED RECOMMENDATION

ศรินทร์ ไสยวัฒนากร

การเพิ่มศักยภาพระบบแนะนำด้วยการผสมผสานการวิเคราะห์ความรู้สึกด้วยแบบจำลอง SieBERT เข้า
กับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้ร่วม



ศรินทร์ ไสภณพัฒน์นกร

การค้นคว้าอิสระนี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร

วิทยาศาสตร์มหาบัณฑิต สาขาวิทยาการข้อมูล

คณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ

ปีการศึกษา 2567

ลิขสิทธิ์ของมหาวิทยาลัยศรีนครินทรวิโรฒ

SENTIMENT ENHANCED CF: UTILIZING SIEBERT FOR OPTIMIZED RECOMMENDATION



KIRIN SOPONWATTANAKORN

An Independent Study Submitted in Partial Fulfillment of the Requirements

for the Degree of MASTER OF SCIENCE

(Data Science)

Faculty of Science, Srinakharinwirot University

2024

Copyright of Srinakharinwirot University

การค้นคว้าอิสระ

เรื่อง

การเพิ่มศักยภาพระบบแนะนำด้วยการผสมผสานการวิเคราะห์ความรู้สึกด้วยแบบจำลอง SieBERT
เข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้งาน

ของ

ศิริน โสภณวัฒนากุล

ได้รับอนุมัติจากบัณฑิตวิทยาลัยให้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร

วิทยาศาสตรมหาบัณฑิต สาขาวิทยาการข้อมูล

ของมหาวิทยาลัยศรีนครินทรวิโรฒ

(รองศาสตราจารย์ นายแพทย์จัตตชัย เอกปัญญากุล)

คณบดีบัณฑิตวิทยาลัย

คณะกรรมการสอบปากเปล่าการค้นคว้าอิสระ

อาจารย์ที่ปรึกษา

ประธานกรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.ศิริสรพ เหล่าหะเกียรติ)

(ผู้ช่วยศาสตราจารย์ ดร.จันตรี ผลประเสริฐ)

กรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.วราภรณ์ วิทยานนท์)

ชื่อเรื่อง	การเพิ่มศักยภาพระบบแนะนำด้วยการผสมผสานการวิเคราะห์ความรู้สึกด้วยแบบจำลอง SieBERT เข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน
ผู้วิจัย	ศรินทร์ โสภณวัฒนากร
ปริญญา	วิทยาศาสตร์มหาบัณฑิต
ปีการศึกษา	2567
อาจารย์ที่ปรึกษา	ผู้ช่วยศาสตราจารย์ ดร.ศิริสรพร เหล่าหะเกียรติ

บทคัดย่อ

การเพิ่มขึ้นอย่างรวดเร็วของข้อมูลรีวิวในอุตสาหกรรมอีคอมเมิร์ซส่งผลให้เกิดปัญหาข้อมูลล้นเกิน (Information Overload) ซึ่งทำให้เกิดความยากต่อการตัดสินใจซื้อสินค้า ระบบแนะนำ (Recommender System) จึงถูกพัฒนาขึ้นเพื่อช่วยคัดกรองข้อมูลและเพิ่มประสิทธิภาพในการตัดสินใจ โดยระบบที่ได้รับความนิยมคือระบบแนะนำแบบการกรองข้อมูลร่วม (Collaborative Filtering) อย่างไรก็ตาม ระบบดังกล่าวยังมีข้อจำกัด เช่น ความกระจัดกระจายของข้อมูล (Sparsity) และการพึ่งพาคะแนนเรตติ้งเพียงอย่างเดียวซึ่งไม่สะท้อนอารมณ์ของผู้บริโภคอย่างแท้จริง งานวิจัยนี้นำเสนออีกหนึ่งแนวทางโดยใช้แบบจำลอง SieBERT สำหรับวิเคราะห์ความรู้สึก (Sentiment Analysis) จากข้อความรีวิวในชุดข้อมูล Amazon Kindle Review แล้วผสมผสานคะแนนความรู้สึก (sentiment scores) เข้ากับคะแนนเรตติ้งด้วยสมการถ่วงน้ำหนัก 3 รูปแบบ จากนั้นนำข้อมูลที่ได้นำไปใช้ในระบบแนะนำ 3 ประเภท ได้แก่ ระบบแนะนำแบบการกรองข้อมูลโดยพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้ (User-based Collaborative Filtering), ระบบแนะนำแบบการกรองข้อมูลโดยพึ่งพาผู้ใช้ร่วมโดยอิงรายการ (Item-based Collaborative Filtering) และระบบแนะนำแบบผสม (Hybrid Collaborative Filtering) ที่รวม User-based และ Item-based เข้าด้วยกันผ่านวิธีการผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) และประเมินประสิทธิภาพด้วยค่า MAE RMSE และ precision@10 ผลการทดลองแสดงให้เห็นว่าการผสมผสานค่า sentiment scores ช่วยเพิ่มความแม่นยำในการคาดการณ์ความสนใจของผู้บริโภคได้อย่างมีนัยสำคัญ โดยเฉพาะอย่างยิ่งการใช้ sentiment scores เพียงอย่างเดียวให้ผลลัพธ์ที่ดีที่สุด ในแง่ของ MAE และ RMSE ในขณะที่การใช้คะแนนเรตติ้งเพียงอย่างเดียวให้ค่า precision@10 สูงที่สุดในทุกระบบแนะนำ

คำสำคัญ: ระบบแนะนำ, การวิเคราะห์ความรู้สึก, การกรองข้อมูลแบบพึ่งพาผู้ใช้งาน, SieBERT

Title	SENTIMENT ENHANCED CF: UTILIZING SIEBERT FOR OPTIMIZED RECOMMENDATION
Author	KIRIN SOPONWATTANAKORN
Degree	MASTER OF SCIENCE
Academic Year	2024
Advisor	Sirisup Laohakiat

The rapid increase in review data within the e-commerce industry has led to the problem of information overload, making purchasing decisions more difficult. Recommender systems have therefore been developed to filter information and enhance decision-making efficiency, with collaborative filtering-based recommender systems being among the most widely adopted in this industry. However, such systems still face limitations, including data sparsity and an overreliance on rating scores, which may not fully capture consumers' true sentiments. This study proposes an alternative approach by utilizing the SieBERT model to perform sentiment analysis on review texts from the Amazon Kindle Review dataset. The resulting sentiment scores are then integrated with the original ratings using three different weighted combination formulas. The enhanced data are subsequently applied to three types of recommender systems: User-based Collaborative Filtering, Item-based Collaborative Filtering, and a Hybrid Collaborative Filtering method that combines User-based and Item-based approaches through a weighted linear combination. System performance is evaluated using MAE, RMSE, and Precision@10 metrics. Experimental results demonstrate that integrating sentiment scores helps significantly improves the prediction accuracy of consumer interests. Notably, using sentiment scores alone yields the best results in terms of MAE and RMSE, while using rating scores alone achieves the highest Precision@10 across all recommender systems.

Keywords: Recommender System Sentiment Analysis Collaborative Filtering SieBERT

กิตติกรรมประกาศ

สารนิพนธ์ฉบับนี้สำเร็จลุล่วงได้ด้วยความช่วยเหลือจาก ผศ.ดร.ศิริสรรพ เหล่าหะเกียรติ อาจารย์ที่ปรึกษา ที่ช่วยเหลือ ให้คำแนะนำในทุกขั้นตอน ทำให้สารนิพนธ์ฉบับนี้เสร็จสมบูรณ์ ข้าพเจ้าขอแสดงความขอบคุณอย่างยิ่งต่อท่านอาจารย์ที่ให้ความเมตตา ชี้แนะแนวทางการวิจัย อย่างเป็นระบบ รวมถึงให้กำลังใจและข้อเสนอแนะที่เป็นประโยชน์ตลอดระยะเวลาการดำเนินงาน นอกจากนี้ ข้าพเจ้าขอขอบคุณคณาจารย์ทุกท่านในสาขาวิชาวิทยาการข้อมูล คณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ ที่ได้ถ่ายทอดความรู้และประสบการณ์อันมีค่า ซึ่งเป็นพื้นฐานสำคัญที่ช่วยเสริมสร้างความเข้าใจในเนื้อหาของงานวิจัย ข้าพเจ้าขอขอบคุณเพื่อน นิสิตปริญญาโททุกคน ที่ได้แลกเปลี่ยนความคิดเห็น และช่วยเหลือในทุกเรื่อง ตลอดจนครอบครัวและเพื่อน ๆ ที่คอยสนับสนุน ให้กำลังใจ และอยู่เคียงข้างในทุกช่วงเวลาของการศึกษาและการจัดทำสารนิพนธ์ฉบับนี้

ศิริณ โสภณวัฒนาก

สารบัญ

หน้า

บทคัดย่อภาษาไทย	ง
บทคัดย่อภาษาอังกฤษ.....	จ
กิตติกรรมประกาศ	ฉ
สารบัญ.....	ช
สารบัญตาราง.....	ณ
สารบัญภาพ.....	ญ
บทที่ 1 บทนำ.....	1
1.1 ความสำคัญและความเป็นมาของงานวิจัย.....	1
1.2 วัตถุประสงค์ของงานวิจัย	2
1.3 ขอบเขตของงานวิจัย.....	2
1.4 ตัวแปรที่ศึกษา	2
1.5 ประโยชน์ที่คาดว่าจะได้รับ	3
บทที่ 2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง	4
2.1 ระบบแนะนำ (Recommender System)	4
2.2 ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน (Collaborative filtering)	9
2.3 ค่าความคล้ายคลึงแบบโคไซน์ (Cosine Similarity)	14
2.4 การวัดประสิทธิภาพของแบบจำลอง (Evaluation Metrics)	15
2.5 การวิเคราะห์ความรู้สึก (Sentiment Analysis)	17
2.6 แบบจำลอง SieBERT	18
2.7 ค่าพารามิเตอร์ (Hyperparameter)	19
2.8 การผสานผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำ	19
2.9 การสร้างระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งานแบบผสม	21
2.10 งานวิจัยที่เกี่ยวข้อง (Related Research)	21
บทที่ 3 การดำเนินงานวิจัย	36
3.1 การประมวลผลข้อมูลเบื้องต้น	37

3.2 การเทรน ปรับแต่ง และประเมินผลแบบจำลอง SieBERT สำหรับการวิเคราะห์	
ความรู้สึก	40
3.3 การสร้างพีเจอร์จาก Sentiment Analysis.....	41
3.4 การนำชุดข้อมูลที่ได้มาใช้กับระบบแนะนำ	42
บทที่ 4 ผลการดำเนินงานวิจัย.....	43
4.1 ผลลัพธ์ของแบบจำลองวิเคราะห์ความรู้สึก	43
4.2 ผลลัพธ์ของการนำคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกไปผสมผสานกับ	
คะแนน การจัดอันดับ(เรตติ้ง) และนำไปใช้ในระบบแนะนำ	44
บทที่ 5 สรุปผลการวิจัย อภิปรายผล และข้อเสนอแนะ.....	58
5.1 สรุปผลการวิจัย	58
5.2 อภิปรายผลการวิจัย.....	59
5.3 ข้อเสนอแนะ	61
บรรณานุกรม	62

สารบัญตาราง

หน้า

ตาราง 1 แสดงผลการทดลองของงานวิจัยที่เกี่ยวข้อง.....	28
ตาราง 2 ชุดข้อมูลที่จะนำไปใช้กับระบบแนะนำ.....	42
ตาราง 3 การปรับปรุงพารามิเตอร์	44
ตาราง 4 การเปรียบเทียบประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 1-5	46
ตาราง 5 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 1-5.....	47
ตาราง 6 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 0-1	48
ตาราง 7 การเปรียบเทียบประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับ (เรตติ้ง) ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 1-5	49
ตาราง 8 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 1-5	50
ตาราง 9 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 0-1	50
ตาราง 10 การเปรียบเทียบประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับของระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนนเรตติ้ง 1-5.....	52

ตาราง 11 สรุปผลลัพธ์ของการผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและ คะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนน เรตติ้ง 1-5.....	56
ตาราง 12 สรุปผลลัพธ์ของการผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและ คะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนน เรตติ้ง 0-1.....	57



สารบัญภาพ

หน้า

ภาพประกอบ 1 Long tail distribution	5
ภาพประกอบ 2 ตัวอย่างระบบแนะนำและผลิตภัณฑ์เป้าหมาย.....	6
ภาพประกอบ 3 ตารางตัวอย่างฐานข้อมูล (user-item matrix) หรือเมทริกซ์ยูทิลิตี้ (Utility matrix)	6
ภาพประกอบ 4 การแบ่งระบบแนะนำตามเทคนิค.....	8
ภาพประกอบ 5 แสดงตัวอย่างวิธีการกรองข้อมูลแบบพึ่งพาผู้ใช้อย่างยิ่ง (User-Based Collaborative Filtering).....	11
ภาพประกอบ 6 ตัวอย่างรูปแบบตารางฐานข้อมูล (user-item matrix)	12
ภาพประกอบ 7 ตารางตัวอย่างการวัดค่าความคล้ายคลึงของการกรองข้อมูลแบบพึ่งพาผู้ใช้อย่างยิ่ง (User-Based Collaborative Filtering) หรือเรียกว่า User-User Matrix โดยมีเป้าหมายเป็นผู้ใช้ที่ 3	12
ภาพประกอบ 8 แสดงตัวอย่างวิธีการกรองข้อมูลแบบพึ่งพาผู้ใช้อย่างยิ่งโดยรายการ (Item-Based Collaborative Filtering)	13
ภาพประกอบ 9 Taxonomy of sentiment analysis technique.....	18
ภาพประกอบ 10 แผนภาพแสดงกระบวนการดำเนินงานของงานวิจัยนี้	36
ภาพประกอบ 11 โครงสร้างข้อมูล Kindle_Store.csv	37
ภาพประกอบ 12 ผลลัพธ์เมื่อเตรียมข้อมูลเรียบร้อยแล้ว	38
ภาพประกอบ 13 การกระจายตัวของข้อมูลในชุด Predict set.....	39
ภาพประกอบ 14 ตัวอย่างชุดข้อมูลที่ทำ sentiment analysis และผลฐานข้อมูลกับค่าเรตติ้งสมการเชิงเส้นถ่วงน้ำหนักแล้ว พร้อมนำไปใช้กับระบบแนะนำ	42
ภาพประกอบ 15 คำสั่งที่ใช้ในการปรับพารามิเตอร์แบบจำลอง SieBERT.....	43
ภาพประกอบ 16 การประเมินประสิทธิภาพของระบบแนะนำ.....	45

บทที่ 1

บทนำ

1.1 ความสำคัญและความเป็นมาของงานวิจัย

ด้วยปริมาณข้อมูลที่เพิ่มขึ้นอย่างต่อเนื่องบนอินเทอร์เน็ต แม้จะมีข้อดีในด้านความหลากหลายของข้อมูล แต่กลับส่งผลให้เกิดภาวะข้อมูลล้นเกิน (Information Overload) ซึ่งทำให้ผู้บริโภคมีความยากลำบากในการตัดสินใจ โดยเฉพาะในอุตสาหกรรมอีคอมเมิร์ซ (e-commerce) ที่มีผู้ค้าและสินค้าจำนวนมาก การเติบโตของการซื้อขายออนไลน์นำไปสู่การให้คะแนนสินค้า (Ratings) และรีวิวจำนวนมาก ซึ่งทำให้ผู้บริโภคเผชิญกับความท้าทายในการตัดสินใจเลือกซื้อ

ระบบแนะนำ (Recommender Systems) ช่วยลดต้นทุนในการค้นหาและเลือกสินค้าในสภาพแวดล้อมการช้อปปิ้งออนไลน์ ในบริบทของอีคอมเมิร์ซ ระบบแนะนำช่วยเพิ่มรายได้จากการขายสินค้าเพิ่มเติม ดังนั้น ความสำคัญของการใช้เทคนิคการแนะนำที่มีประสิทธิภาพและแม่นยำจึงไม่สามารถมองข้ามได้

เนื่องจากมนุษย์มีความสามารถในการตัดสินใจในแต่ละวันอย่างจำกัด ระบบแนะนำจึงมีความสำคัญอย่างยิ่งในการลดเวลาและเพิ่มคุณภาพในการตัดสินใจของผู้ซื้อ โดยการแนะนำสินค้าที่คาดว่าจะตรงกับความต้องการ ซึ่งหากระบบสามารถทำงานได้อย่างแม่นยำ จะทำให้การตัดสินใจซื้อรวดเร็วและง่ายขึ้น

ทั้งนี้ ระบบแนะนำในปัจจุบันสามารถแบ่งออกเป็นหลายประเภท ในการวิจัยนี้ จะเลือกใช้ระบบแนะนำที่ใช้เทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) เนื่องจากเป็นระบบที่แพร่หลายในอุตสาหกรรมอีคอมเมิร์ซ โดยระบบนี้อิงจากข้อมูลผู้ใช้ (user), รายการสินค้า (item), และเรตติ้ง (rating) เท่านั้น ทำให้สามารถนำไปปรับใช้ในธุรกิจอีคอมเมิร์ซได้ทุกขนาด

อย่างไรก็ตาม การใช้เพียงข้อมูลเหล่านี้มีข้อจำกัด เนื่องจากคะแนนเรตติ้งเพียงอย่างเดียวไม่สามารถสะท้อนความคิดเห็น อารมณ์ หรือพฤติกรรมของผู้บริโภคได้อย่างแท้จริง นอกจากนี้ ผู้บริโภคแต่ละคนมีเกณฑ์การให้คะแนนที่แตกต่างกัน ซึ่งอาจนำไปสู่ความเข้าใจที่คลาดเคลื่อนในความหมายของคะแนน เช่น คะแนน 3 อาจหมายถึงไม่ประทับใจหรือพอใช้ก็ได้ อีกทั้งระบบแนะนำแบบ Collaborative Filtering ยังเผชิญกับปัญหาความกระจัดกระจาย (Sparsity) หรือข้อมูลว่างใน utility matrix ซึ่งส่งผลให้การแนะนำมีความไม่แม่นยำ

ในปัจจุบัน การรีวิวสินค้าออนไลน์ได้รับความนิยม ข้อมูลรีวิวมียอดดีคือสามารถนำมาวิเคราะห์เพื่อทำความเข้าใจความคิดเห็นและพฤติกรรมของผู้บริโภค อีกทั้งยังสามารถแปลงเป็น

ค่า sentiment scores และนำไปใช้ร่วมกับระบบแนะนำ เพื่อลดปัญหาความกระจัดกระจายและเพิ่มประสิทธิภาพของระบบ

ในการวิจัยนี้ ผู้วิจัยได้นำข้อมูลรีวิวสินค้ามาวิเคราะห์อารมณ์ (Sentiment Analysis) โดยใช้แบบจำลอง Fine-Tuned SieBERT ซึ่งเป็นแบบจำลองที่ปรับแต่งจาก RoBERTa สำหรับการวิเคราะห์อารมณ์โดยเฉพาะ หลังจากปรับค่าพารามิเตอร์ของ SieBERT แล้ว ผู้วิจัยได้ทำการรวมค่า sentiment scores กับเรตติ้ง โดยให้น้ำหนักต่างกันระหว่างค่า sentiment scores และค่า ratings ได้แก่ 0, 0.1, 0.3, 0.5, 0.7, 0.9 และ 1

1.2 วัตถุประสงค์ของงานวิจัย

1. เพื่อนำค่า sentiment scores จากการทำ sentiment analysis มาเพิ่มประสิทธิภาพให้กับระบบแนะนำแบบ Collaborative filtering
2. เพื่อศึกษาและเปรียบเทียบค่าน้ำหนักระหว่าง sentiment scores และ rating ที่เพิ่มประสิทธิภาพของระบบแนะนำได้มากที่สุด

1.3 ขอบเขตของงานวิจัย

งานวิจัยนี้ใช้ข้อมูลรีวิวสินค้า Amazon Kindle Review จากแหล่งข้อมูลสาธารณะ โดยมีเป้าหมายคือการนำข้อมูลรีวิวมาทำ sentiment analysis เพื่อนำ sentiment label ที่ได้ไปเพิ่มประสิทธิภาพและลดข้อจำกัดให้กับระบบแนะนำ โดยแบบจำลองที่เลือกใช้ทำ sentiment analysis คือ SieBERT และระบบแนะนำที่เลือกนำมาทดสอบประสิทธิภาพคือ User-based Collaborative filtering, Item-based Collaborative filtering และ Hybrid Collaborative filtering โดยวัดผลของแบบจำลอง sentiment analysis ด้วย Accuracy, AUC, F1-Score และวัดประสิทธิภาพของระบบแนะนำด้วย MAE RMSE และ Precision@K

1.4 ตัวแปรที่ศึกษา

1. user_id หรือรหัสของผู้เขียนรีวิว
2. parent_asin หรือรหัสของสินค้าหลัก
3. title หรือชื่อเรื่องของรีวิวที่ผู้ใช้เขียน
4. text หรือเนื้อหาของรีวิวที่ผู้ใช้เขียน
5. rating หรือคะแนนของสินค้าที่ได้รับการรีวิว (ตั้งแต่ 1.0 ถึง 5.0)

1.5 ประโยชน์ที่คาดว่าจะได้รับ

1. สามารถนำค่า sentiment scores จากการทำ sentiment analysis มาเพิ่มประสิทธิภาพให้กับระบบแนะนำแบบ Collaborative filtering
2. สามารถเปรียบเทียบค่าน้ำหนักระหว่าง sentiment scores และ rating ที่เพิ่มประสิทธิภาพของระบบแนะนำได้มากที่สุด



บทที่ 2

ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

ในงานวิจัยนี้ ผู้วิจัยได้ศึกษาเกี่ยวกับการปรับปรุงระบบแนะนำแบบ Collaborative Filtering ให้มีประสิทธิภาพมากขึ้นและลดข้อจำกัดลงด้วย sentiment score จากการทำ sentiment analysis จากข้อมูลรีวิว โดยเสนอหัวข้อดังต่อไปนี้

1. ระบบแนะนำ (Recommender Systems)
2. ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน (Collaborative filtering)
3. ค่าความคล้ายคลึงแบบโคไซน์ (Cosine Similarity)
4. การวัดประสิทธิภาพของแบบจำลอง (Evaluation Metrics)
5. การวิเคราะห์ความรู้สึก (Sentiment Analysis)
6. แบบจำลอง SieBERT
7. ค่าพารามิเตอร์ (Hyperparameter)
8. การผสานผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำ
9. งานวิจัยที่เกี่ยวข้อง (Related Research)

2.1 ระบบแนะนำ (Recommender System)

ระบบแนะนำ (Recommender System) คือระบบกรองข้อมูลที่ช่วยแก้ปัญหาการมีข้อมูลล้นเกิน โดยทำหน้าที่คัดกรองข้อมูลที่สำคัญจากข้อมูลจำนวนมากที่เกิดขึ้นอย่างต่อเนื่องตามความสนใจหรือพฤติกรรมการใช้งานของผู้ใช้ที่เกี่ยวข้องกับสินค้า ระบบแนะนำมีความสามารถในการทำนายว่าผู้ใช้จะสนใจสินค้าใดสินค้าหนึ่งหรือไม่ โดยอิงจากข้อมูลของผู้ใช้ (Isinkaye, Folajimi, & Ojokoh, 2015)

ระบบแนะนำมีฐานรากมาจากแนวคิด "Long Tail" ที่ Chris Anderson บรรณาธิการนิตยสาร Wired ได้เสนอขึ้นเมื่อปี ค.ศ. 2006 แนวคิด "Long Tail" คือแนวคิดที่แบ่งสินค้าออกเป็นสองส่วน คือ "หัว" (Head) ซึ่งหมายถึงสินค้าที่ขายดี และ "หาง" (Tail) คือสินค้าที่ขายได้น้อยกว่า แต่เมื่อนำสินค้าจำนวนมากในส่วนหางมารวมกัน กลับมียอดขายรวมมากกว่าส่วนหัวเสียอีก ระบบแนะนำจึงเน้นการแนะนำสินค้าส่วนนี้ให้เจาะลงไปทีละลูกค้าแต่ละรายมากขึ้น (Patcharawongsakda, 2021)

The Long Tail



ภาพประกอบ 1 Long tail distribution

ที่มา: <https://www.nngroup.com/articles/long-tail/>(Sunwall, 2021)

ทั้งนี้ ระบบแนะนำได้รับการใช้งานอย่างแพร่หลาย และมีตัวอย่างที่ชัดเจนมากมาย

ดังภาพประกอบ 2

ระบบแนะนำ	ผลิตภัณฑ์เป้าหมาย
Amazon.com	หนังสือและผลิตภัณฑ์อื่นๆ
Netflix	ทีวีดี, วิดีโอสตรีมมิ่ง
GroupLens	ข่าว
MovieLens	ภาพยนตร์
last.fm	เพลง
Google News	ข่าว
Google Search	โฆษณา
Facebook	เพื่อน, โฆษณา
Pandora	เพลง
YouTube	วิดีโอออนไลน์
Tripadvisor	ผลิตภัณฑ์เกี่ยวกับการท่องเที่ยว
IMDb	ภาพยนตร์

ภาพประกอบ 2 ตัวอย่างระบบแนะนำและผลิตภัณฑ์เป้าหมาย

ที่มา: <https://rd.springer.com/book/10.1007/978-3-319-29659-3>

ในกระบวนการสร้างระบบแนะนำ ส่วนสำคัญที่ไม่สามารถขาดได้ในระบบแนะนำคือ X = เซตของลูกค้าหรือผู้ใช้ (customers/users) และ S = เซตของรายการ (items) ซึ่งจะถูกรวมเข้าด้วยกันในฟังก์ชันยูทิลิตี้ $u: X \times S \rightarrow R$ เพื่อสร้างเมทริกซ์ยูทิลิตี้ (Utility matrix) ดังภาพประกอบ 3

ผู้ใช้ / รายการ	รายการ A	รายการ B	รายการ C	รายการ D	รายการ E
ผู้ใช้ 1	5	-	-	1	-
ผู้ใช้ 2	4	-	4	-	2
ผู้ใช้ 3	-	-	5	-	-
ผู้ใช้ 4	1	-	-	5	-
ผู้ใช้ 5	-	3	-	-	5

ภาพประกอบ 3 ตารางตัวอย่างฐานข้อมูล (user-item matrix) หรือเมทริกซ์ยูทิลิตี้ (Utility matrix)

ในตารางตัวอย่างนี้ เลข 1-5 แทนคะแนนที่ผู้ใช้ให้กับแต่ละรายการ(เรตติ้ง) ส่วน "-" แทนการที่ผู้ใช้ยังไม่ได้ให้คะแนนรายการนั้น ๆ

สามารถกล่าวได้ว่าตารางตัวอย่างฐานข้อมูล (user-item matrix) ก็คือเมทริกซ์ยูทิลิตี้ (Utility matrix) ที่ยังไม่สมบูรณ์ โดยข้อมูลที่ถูกลบออกไปนั้นมีจำนวนน้อยมาก จุดมุ่งหมายของระบบแนะนำคือการพยายามทำนายเรตติ้งในช่องว่างเหล่านี้ด้วยวิธีการต่าง ๆ โดยไม่จำเป็นต้องเติมให้เต็ม 100% เนื่องจากจะสิ้นเปลืองทรัพยากรมาก โดยจะเลือกเติมข้อมูลเฉพาะในช่องที่มีแนวโน้มเรตติ้งสูง (high rating) เท่านั้น ทำให้สามารถให้การแนะนำแก่ผู้ใช้ (user) มีประสิทธิภาพ และช่วยระบุว่าผู้ใช้น่าจะชอบรายการใดบ้าง

ในเมทริกซ์ยูทิลิตี้ (Utility matrix) จะมีการเติมข้อมูลที่เรียกว่าเรตติ้ง (Ratings) ซึ่งแสดงถึงความชอบของผู้ใช้ต่อรายการต่าง ๆ โดยเรตติ้งสามารถแบ่งออกเป็นสองประเภท คือ

1. เรตติ้งแบบชัดเจน (Explicit rating) หมายถึง ระบบจะกระตุ้นให้ผู้ใช้ให้คะแนนหรือขอเสนอแนะสำหรับรายการต่าง ๆ ผ่านทางอินเตอร์เฟซของระบบ ซึ่งต้องการความพยายามจากผู้ใช้ ทำให้การเก็บข้อมูลทำได้ยาก แต่มีความแม่นยำสูง เช่น เลขเรตติ้งหรือข้อมูลรีวิวสินค้าที่ชัดเจน ซึ่งในงานวิจัยของเราจะทำงานกับเรตติ้งชนิดนี้เป็นหลัก

2. เรตติ้งแบบแฝง (Implicit rating) หมายถึง อาศัยการอนุมานความชอบของผู้ใช้จากพฤติกรรม เช่น การดูภาพยนตร์จนจบ หรือเวลาในการใช้งานเว็บเพจ ซึ่งไม่ต้องการความพยายามจากผู้ใช้ในการให้ข้อมูล ทำให้เก็บข้อมูลได้ง่ายกว่า แม้จะความแม่นยำจะต่ำกว่า (อาทิ การดูภาพยนตร์จนจบไม่ได้แปลว่าเราชอบภาพยนตร์เรื่องนั้น เป็นต้น)

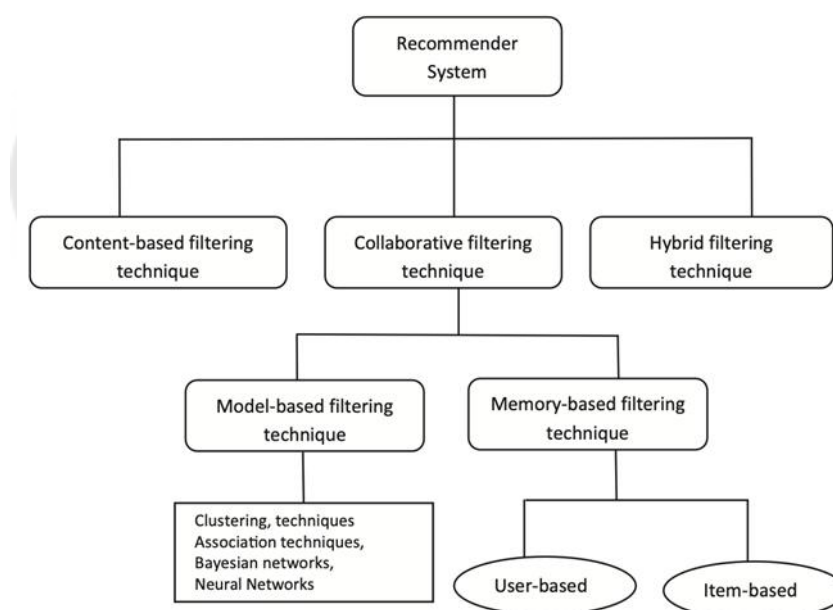
ปัจจุบัน ระบบแนะนำพื้นฐานแบ่งออกเป็นสามแนวทางใหญ่ ๆ ได้แก่

1. ระบบแนะนำตามเนื้อหา (Content-based Filtering) ระบบนี้ใช้การวิเคราะห์คุณสมบัติของรายการ (Item) เพื่อนำเสนอคำแนะนำ โดยพิจารณาจากโปรไฟล์ของผู้ใช้ที่มีข้อมูลเนื้อหาของรายการที่ผู้ใช้เคยประเมินมาก่อน ระบบจะแนะนำรายการที่คล้ายกับรายการที่ผู้ใช้เคยให้คะแนนในเชิงบวก โดยไม่จำเป็นต้องอ้างอิงโปรไฟล์ของผู้ใช้คนอื่น อย่างไรก็ตาม ระบบนี้จำเป็นต้องมีข้อมูลคุณลักษณะของรายการและโปรไฟล์ผู้ใช้ที่ละเอียด และมักพบข้อจำกัดในเรื่องของการจำกัดเนื้อหา (Content Overspecialization) ซึ่งคำแนะนำจะอยู่ในกรอบของรายการที่คล้ายคลึงกับสิ่งที่ผู้ใช้เคยชอบเท่านั้น

2. ระบบแนะนำโดยใช้การกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) ระบบนี้อาศัยข้อมูลการปฏิสัมพันธ์ระหว่างผู้ใช้กับรายการ เช่น การให้เรตติ้ง เพื่อวิเคราะห์ความคล้ายคลึง โดยใช้การกรองตามผู้ใช้ เช่น หากผู้ใช้ n คนให้เรตติ้งรายการที่คล้ายกันแต่ผู้ใช้

เป้าหมายยังไม่เคยให้เรตติ้ง ระบบอาจแนะนำรายการที่ผู้ใช้ n คน ให้เรตติ้งสูงแก่ผู้ใช้เป้าหมาย เนื่องจากงานวิจัยนี้มุ่งเน้นระบบแนะนำประเภทนี้เป็นหลัก รายละเอียดเพิ่มเติมจะถูกอธิบายต่อไปในหัวข้อที่ 2

3. ระบบแนะนำแบบผสม (Hybrid Recommender Systems) ระบบนี้ใช้การรวมเทคนิคหลายรูปแบบ เพื่อเพิ่มความแม่นยำและลดข้อจำกัดของระบบแนะนำเดี่ยว แนวคิดหลักคือ การผสมอัลกอริทึมหลายชนิดเข้าด้วยกัน เช่น การผสมการกรองตามเนื้อหา (Content-Based Filtering) กับการกรองข้อมูลแบบพึ่งพาผู้ใช้อย่างร่วม (Collaborative Filtering) หรือการรวมการกรองข้อมูลแบบพึ่งพาผู้ใช้อย่างร่วมตามผู้ใช้ (User-Based) กับแบบรายการ (Item-Based) ซึ่งการผสมนี้ช่วยให้ระบบแนะนำมีประสิทธิภาพมากขึ้น เนื่องจากจุดอ่อนของอัลกอริทึมหนึ่งจะถูกชดเชยโดยอีกอัลกอริทึมหนึ่งได้



ภาพประกอบ 4 การแบ่งระบบแนะนำตามเทคนิค

ที่มา: [https://www.sciencedirect.com/science/article/pii/S1110866515000341?](https://www.sciencedirect.com/science/article/pii/S1110866515000341?via%3Dihub)

via%3Dihub

จากภาพประกอบ 4 แสดงการแบ่งระบบแนะนำพื้นฐานออกเป็นสามแนวทางใหญ่ ๆ ได้แก่

1. ระบบแนะนำตามเนื้อหา (Content-based Filtering) ระบบนี้ใช้การวิเคราะห์คุณสมบัติของรายการ (Item) เพื่อนำเสนอคำแนะนำ โดยพิจารณาจากโปรไฟล์ของผู้ใช้ที่มีข้อมูลเนื้อหาของรายการที่ผู้ใช้เคยประเมินมาก่อน ระบบจะแนะนำรายการที่คล้ายกับรายการที่ผู้ใช้เคยให้คะแนนในเชิงบวก โดยไม่จำเป็นต้องอ้างอิงโปรไฟล์ของผู้ใช้คนอื่น อย่างไรก็ตาม ระบบนี้จำเป็นต้องมีข้อมูลคุณลักษณะของรายการและโปรไฟล์ผู้ใช้ที่ละเอียด และมักพบข้อจำกัดในเรื่องของการจำกัดเนื้อหา (Content Overspecialization) ซึ่งคำแนะนำจะอยู่ในกรอบของรายการที่คล้ายคลึงกับสิ่งที่ผู้ใช้เคยชอบเท่านั้น

2. ระบบแนะนำโดยใช้การกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) ระบบแนะนำนี้อาศัยข้อมูลการปฏิสัมพันธ์ระหว่างผู้ใช้กับรายการ เช่น การให้เรตติ้ง เพื่อวิเคราะห์ความคล้ายคลึง โดยใช้การกรองตามผู้ใช้ เช่น หากผู้ใช้ n คนให้เรตติ้งรายการที่คล้ายกันแต่ผู้ใช้เป้าหมายยังไม่เคยให้เรตติ้ง ระบบอาจแนะนำรายการที่ผู้ใช้ n คน ให้เรตติ้งสูงแก่ผู้ใช้เป้าหมาย เนื่องจากงานวิจัยนี้มุ่งเน้นระบบแนะนำประเภทนี้เป็นหลัก รายละเอียดเพิ่มเติมจะถูกอธิบายต่อไปในหัวข้อที่ 2

3. ระบบแนะนำแบบผสม (Hybrid Recommender Systems) ระบบนี้ใช้การรวมเทคนิคหลายรูปแบบ เพื่อเพิ่มความแม่นยำและลดข้อจำกัดของระบบแนะนำเดี่ยว แนวคิดหลักคือการผสมอัลกอริทึมหลายชนิดเข้าด้วยกัน เช่น การผสมการกรองตามเนื้อหา (Content-Based Filtering) กับการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) หรือการรวมการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมตามผู้ใช้ (User-Based) กับแบบรายการ (Item-Based) ซึ่งการผสมนี้ช่วยให้ระบบแนะนำมีประสิทธิภาพมากขึ้น เนื่องจากจุดอ่อนของอัลกอริทึมหนึ่งจะถูกชดเชยโดยอีกอัลกอริทึมหนึ่งได้

งานวิจัยนี้เลือกใช้ทั้งระบบแนะนำแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) และระบบแนะนำแบบผสม (Weighted Hybrid Recommender Systems) จึงจะเน้นรายละเอียดในสองส่วนนี้เป็นพิเศษ

2.2 ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative filtering)

คำแนะนำจากระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมสามารถแบ่งออกเป็นสองประเภท คือ การคาดการณ์ค่าคะแนนของการรวมกันระหว่างผู้ใช้กับสินค้า (Predicting the rating value of a user-item combination) และการกำหนดสินค้าหรือผู้ใช้ที่เกี่ยวข้องที่สุด

(Determining the top-k items or top-k users) โดยในงานวิจัยนี้จะโฟกัสไปที่การคาดการณ์ค่าคะแนนของการรวมกันระหว่างผู้ใช้กับสินค้าเป็นหลัก

ระบบแนะนำที่ใช้เทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน (Collaborative filtering) เริ่มต้นด้วยการสร้างฐานข้อมูล (user-item matrix) ซึ่งบันทึกความชอบ (เรตติ้ง) ของผู้ใช้ต่อรายการต่าง ๆ จากนั้นจะคำนวณค่าความคล้ายคลึงระหว่างโปรไฟล์ผู้ใช้หรือรายการ เพื่อแนะนำรายการที่เหมาะสม โดยกลุ่มผู้ใช้ที่มีความสนใจใกล้เคียงกันจะถูกจัดเป็น "เพื่อนบ้าน" (Neighborhood) ซึ่งทำให้ถูกมีอีกชื่อเรียกว่า Neighborhood-Based Collaborative filtering ระบบจะเสนอรายการที่ผู้ใช้ยังไม่เคยให้คะแนน แต่ได้รับเรตติ้งในเชิงบวกจากเพื่อนบ้านที่มีโปรไฟล์คล้ายกัน

อย่างไรก็ตาม ระบบแนะนำที่ใช้เทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน (Collaborative filtering) มีข้อจำกัดอยู่พอสมควร หนึ่งในนั้นคือปัญหา "cold-start" ซึ่งหมายถึงสถานการณ์ที่ระบบไม่สามารถแนะนำรายการได้เมื่อเพื่อนบ้านไม่มีเรตติ้งเลย และนอกจากนี้ยังมีข้อจำกัดสำคัญที่งานวิจัยของเราสนใจคือ "sparsity" ซึ่งหมายถึงการมีช่องว่างในตารางมากเกินไป ทำให้ไม่สามารถระบุรูปแบบ (pattern) ที่ชัดเจนได้ อย่างไรก็ตาม ปัญหา "sparsity" สามารถแก้ไขได้หลายวิธี เช่น ระบบแนะนำแบบผสม (hybrid approach) ซึ่งสามารถจัดการทั้งปัญหา cold-start และ sparsity รวมถึงวิธีการผสมผสาน sentiment score เข้ากับระบบแนะนำ ซึ่งเป็นแนวทางที่น่าสนใจในการเพิ่มประสิทธิภาพของระบบแนะนำ และเป็นแนวทางที่งานวิจัยนี้เลือกใช้

ในเชิงเทคนิค ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน แบ่งออกได้เป็นสองรูปแบบหลักคือ แบบใช้หน่วยความจำ (Memory-Based) และแบบโมเดล (Model-Based)

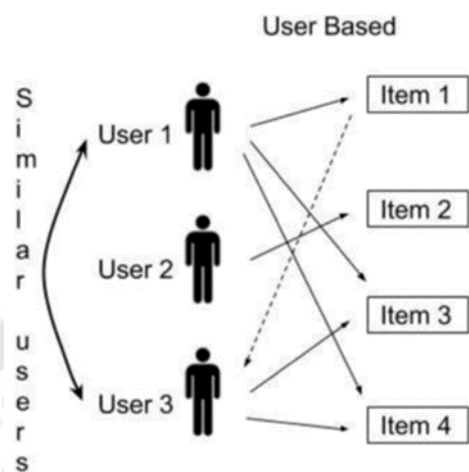
2.2.1 ระบบแนะนำแบบกรองข้อมูลพึ่งพาผู้ใช้งานโดยวิธีอิงหน่วยความจำ (Memory-Based Collaborative Filtering Recommendation System)

ระบบแนะนำโดยวิธีกรองข้อมูลพึ่งพาผู้ใช้งานแบบอิงหน่วยความจำ (Memory-Based Collaborative Filtering Recommendation System) เป็นวิธีแรก ๆ ที่พัฒนาขึ้นสำหรับการกรองข้อมูลแบบพึ่งพาผู้ใช้งานที่ยังคงใช้กันอยู่อย่างแพร่หลายในปัจจุบัน วิธีนี้อิงจากข้อเท็จจริงที่ว่าผู้ใช้ที่คล้ายกันมักจะแสดงรูปแบบพฤติกรรมการให้คะแนนที่คล้ายกัน ซึ่งสามารถแบ่งออกได้เป็น 2 ประเภท คือ

1) การกรองข้อมูลแบบพึ่งพาผู้ใช้งานโดยอิงผู้ใช้ (User-Based Collaborative Filtering)

วิธีการกรองข้อมูลแบบพึ่งพาผู้ใช้งานโดยอิงผู้ใช้ (User-Based Collaborative Filtering) เป็นเทคนิคที่ใช้ข้อมูลจากผู้ใช้ (User) ที่มีความชอบใกล้เคียงกับผู้ใช้เป้าหมาย (Target User) ในการสร้างคำแนะนำ หลักการคือการค้นหาผู้ใช้กลุ่มที่มีรสนิยมใกล้เคียงกับผู้ใช้เป้าหมาย

แล้วนำค่าเฉลี่ยถ่วงน้ำหนักของเรตติ้ง (Weighted Average of Ratings) จากกลุ่มนี้มาทำนายเรตติ้งของผู้ใช้เป้าหมายที่ยังขาดอยู่



ภาพประกอบ 5 แสดงตัวอย่างวิธีการกรองข้อมูลแบบพึ่งพาผู้ใช้รวมโดยอิงผู้ใช้ (User-Based Collaborative Filtering)

ที่มา: https://www.researchgate.net/figure/User-based-and-Item-based-Collaborative-Filtering_fig2_366902172

ตัวอย่างเช่น หากผู้ใช้ชื่อ "ปิติ" และ "ชูใจ" มีรูปแบบการให้เรตติ้งภาพยนตร์ที่คล้ายกัน หาก "ปิติ" ให้เรตติ้งสูงแก่ภาพยนตร์เรื่อง Terminator เราสามารถนำข้อมูลนี้มาทำนายว่า "ชูใจ" อาจชอบภาพยนตร์เรื่องนี้เช่นกัน การคาดการณ์มักใช้ข้อมูลจากผู้ใช้ที่คล้ายกับผู้ใช้เป้าหมายมากที่สุดจำนวน k คน โดยคำนวณค่าความคล้ายคลึง (Similarity) ระหว่างแถวของเมทริกซ์เรตติ้ง (Ratings Matrix) เพื่อระบุผู้ใช้ที่มีรสนิยมใกล้เคียงกัน ทั้งนี้ วิธีการวัดความคล้ายคลึงระหว่างผู้ใช้ เช่น การใช้ Cosine Similarity (รายละเอียดเพิ่มเติมในหัวข้อที่ 3) นั้นสามารถช่วยเพิ่มความแม่นยำในการระบุผู้ใช้ที่มีความชอบคล้ายกัน

ผู้ใช้ / ภาพยนตร์	Terminator	Titanic	Inception	Matrix	Avatar
ปิติ	5	3	-	4	-
ซูใจ (ผู้ใช้เป้าหมาย)	-	3	5	-	4
สร้อย	4	-	4	5	-
ออม	-	4	3	-	5
บอย	5	3	4	2	-

ภาพประกอบ 6 ตัวอย่างรูปแบบตารางฐานข้อมูล (user-item matrix)

Item-Id ⇒	1	2	3	4	5	6	Mean Rating	Cosine($i, 3$) (user-user)
User-Id ↓								
1	7	6	7	4	5	4	5.5	0.956
2	6	7	?	4	3	4	4.8	0.981
3	?	3	3	1	1	?	2	1.0
4	1	2	2	3	3	4	2.5	0.789
5	1	?	1	2	3	3	2	0.645

ภาพประกอบ 7 ตารางตัวอย่างการวัดค่าความคล้ายคลึงของการกรองข้อมูลแบบพึ่งพาผู้ใช้งานโดยอิงผู้ใช้ (User-Based Collaborative Filtering) หรือเรียกว่า User-User Matrix โดยมีเป้าหมายเป็นผู้ใช้ที่ 3

ที่มา: <https://rd.springer.com/book/10.1007/978-3-319-29659-3>

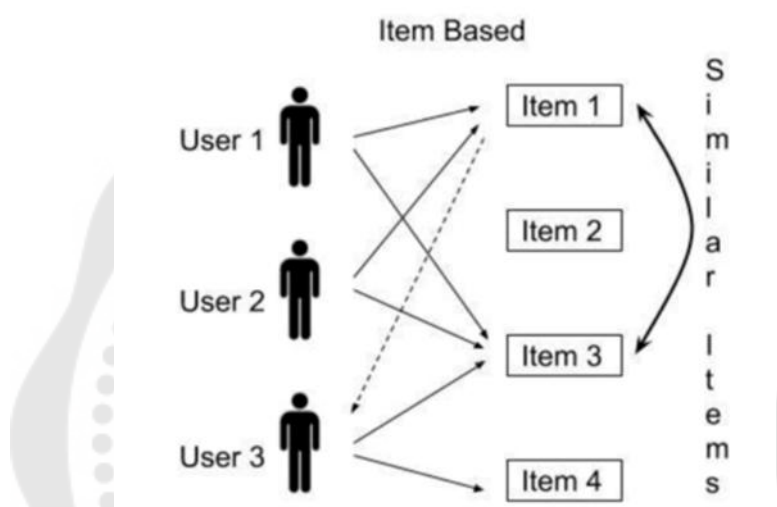
จากภาพประกอบ 7 จะเห็นว่าผู้ใช้ที่มีค่าความคล้ายคลึงใกล้เคียงกับผู้ใช้ที่ 3 ซึ่งเป็นผู้ใช้เป้าหมายคือ ผู้ใช้ที่ 2 ดังนั้น รายการ (Item) ที่จะถูกแนะนำให้ผู้ใช้งานที่ 3 คือ รายการหมายเลข 1 (สามารถดูรายละเอียดเพิ่มเติมเกี่ยวกับค่าความคล้ายคลึงได้ในหัวข้อที่ 3)

ข้อดีของเทคนิคนี้คือสามารถนำเสนอรายการที่หลากหลายและมีโอกาสแนะนำรายการที่ผู้ใช้งานไม่เคยเห็นหรือพิจารณามาก่อน อย่างไรก็ตาม วิธีนี้มีข้อจำกัดเรื่องปัญหาความเบาบางของข้อมูล (Data Sparsity) และปัญหา cold-start เมื่อระบบยังมีข้อมูลจากผู้ใช้งานไม่เพียงพอ เช่น เมื่อผู้ใช้งานที่ยังไม่มีเรตติ้งในระบบเข้ามา

โดยสรุป User-Based Collaborative Filtering เป็นเทคนิคที่เหมาะสมสำหรับระบบแนะนำในสภาพแวดล้อมที่ผู้ใช้งานมีพฤติกรรมการให้เรตติ้งที่ชัดเจนและมีปริมาณเพียงพอ ซึ่งจะช่วยให้สามารถค้นหากลุ่มผู้ใช้ที่มีความชอบใกล้เคียงได้อย่างมีประสิทธิภาพ

2) การกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงรายการ (Item-Based Collaborative Filtering)

วิธีการกรองข้อมูลโดยอิงรายการ (Item-Based Collaborative Filtering) ใช้ข้อมูลจากรายการ (Item) ที่มีความคล้ายคลึงกันเพื่อสร้างคำแนะนำให้แก่ผู้ใช้เป้าหมาย (Target User) หลักการคือการค้นหารายการที่มีความคล้ายคลึงกับรายการเป้าหมาย (Target Item) แล้วใช้เรตติ้งจากผู้ใช้ที่เคยให้กับรายการที่คล้ายคลึงกันในการทำนายเรตติ้งที่ยังขาดอยู่



ภาพประกอบ 8 แสดงตัวอย่างวิธีการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงรายการ (Item-Based Collaborative Filtering)

ที่มา: https://www.researchgate.net/figure/User-based-and-Item-based-Collaborative-Filtering_fig2_366902172

ตัวอย่างเช่น หากเราต้องการทำนายเรตติ้งที่ "ซูโจ" อาจให้กับภาพยนตร์เรื่อง Terminator ระบบจะค้นหารายการอื่นที่มีความคล้ายกับ Terminator เช่น Inception และ The Matrix โดยคำนวณค่าความคล้ายคลึงระหว่างรายการในเมทริกซ์ผู้ใช้-รายการ (User-Item Matrix) จากนั้นนำเรตติ้งของซูโจที่เคยให้กับ Inception และ The Matrix มาช่วยทำนายความเป็นไปได้ที่ซูโจจะชื่นชอบ Terminator เช่นกัน

การวัดค่าความคล้ายคลึงระหว่างรายการมักใช้เทคนิค เช่น Cosine Similarity หรือ Pearson Correlation เพื่อวัดมุมหรือแนวโน้มของเรตติ้งของรายการในมิติเดียวกัน

ยิ่งค่าความคล้ายคลึงสูง รายการจะยิ่งมีแนวโน้มที่จะมีผู้ใช้ให้คะแนนใกล้เคียงกัน ทำให้การคาดการณ์แม่นยำขึ้น

ข้อดีของ Item-Based Collaborative Filtering คือมีความคงที่ เนื่องจากข้อมูลที่เปลี่ยนแปลงมักเกี่ยวข้องกับการเพิ่มหรือลดเรตติ้งของรายการซึ่งไม่ต้องคำนวณความคล้ายคลึงระหว่างรายการใหม่ทุกครั้งเมื่อมีผู้ใช้ใหม่ ข้อจำกัดคือยังคงมีปัญหา cold-start เมื่อมีรายการใหม่เข้ามาในระบบที่ยังไม่มีเรตติ้งจากผู้ใช้ใดเลย ส่งผลให้การแนะนำทำได้ยากขึ้น

โดยสรุป Item-Based Collaborative Filtering มีประโยชน์เมื่อมีรายการที่ได้รับเรตติ้งจากผู้ใช้จำนวนมาก ทำให้สามารถใช้รายการที่คล้ายกันในการสร้างคำแนะนำที่แม่นยำและน่าสนใจให้กับผู้ใช้

2.2.2 ระบบแนะนำแบบกรองข้อมูลพึ่งพาผู้ใช้ร่วมโดยวิธีอิงแบบจำลอง (Model-Based Collaborative Filtering Recommendation System)

ระบบแนะนำแบบกรองข้อมูลพึ่งพาผู้ใช้ร่วมโดยใช้แบบจำลอง (Model-Based Collaborative Filtering Recommendation System) คือการใช้แบบจำลองเพื่อเพิ่มประสิทธิภาพของระบบแนะนำประเภทกรองข้อมูลพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) การสร้างระบบแนะนำวิธีนี้ไม่เพียงแต่ใช้การเรียนรู้ของเครื่อง (Machine Learning) แต่ยังใช้การทำเหมืองข้อมูล (Data Mining) ด้วย เทคนิคเหล่านี้สามารถแนะนำรายการได้รวดเร็ว ตัวอย่างเช่น เทคนิคการลดมิติ (Dimensionality Reduction) เช่น Singular Value Decomposition (SVD), เทคนิคการเติมตาราง (Matrix Completion), วิธีการ Latent Semantic, รวมถึงการถดถอย (Regression) และการจัดกลุ่ม (Clustering)

การใช้อัลกอริทึมการเรียนรู้ (Learning Algorithms) ยังช่วยเปลี่ยนวิธีการแนะนำจากการแนะนำสิ่งที่ผู้ใช้ควรบริโภค ไปสู่การแนะนำเวลาที่เหมาะสมในการบริโภคสินค้า จึงจำเป็นต้องศึกษาอัลกอริทึมการเรียนรู้อื่น ๆ ที่ใช้ในระบบแนะนำแบบใช้แบบจำลอง (Model-based Recommender Systems)

2.3 ค่าความคล้ายคลึงแบบโคไซน์ (Cosine Similarity)

Cosine Similarity เป็นเทคนิคที่ใช้วัดความคล้ายคลึงระหว่างข้อมูลสองชุด โดยข้อมูลแต่ละชุดจะแสดงในรูปของเวกเตอร์ ค่าความคล้ายคลึงนี้พิจารณาจากมุมระหว่างเวกเตอร์ ซึ่งใช้บอกว่าเวกเตอร์ทั้งสองมีแนวโน้มในทิศทางเดียวกันหรือไม่ การคำนวณ Cosine Similarity นั้นหาจากการทำ dot product ระหว่างเวกเตอร์สองชุด แล้วหารด้วยผลคูณของขนาด (magnitude) ของแต่ละเวกเตอร์ U และ V ตามสมการ

$$s(\vec{u}, \vec{v}) = \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| * |\vec{v}|} = \frac{\sum_i r_{u,i} r_{v,i}}{\sqrt{\sum_i r_{u,i}^2} * \sqrt{\sum_i r_{v,i}^2}}$$

U = รายการของชุด U

V = รายการของชุด V

N = จำนวนคู่ของรายการชุด A และรายการชุด B ทั้งหมด

เมื่อนำมาใช้ในระบบแนะนำ (Recommender System) เราสามารถมองว่า user profile และ item profile เป็นเวกเตอร์ใน space เดียวกัน การวัดระยะห่างระหว่างเวกเตอร์เหล่านี้จะช่วยบ่งบอกถึงระดับความคล้ายคลึงของผู้ใช้สินค้าต่าง ๆ ค่าที่น้อยจะแสดงถึงองศาของมุมที่น้อย ซึ่งหมายถึงความคล้ายคลึงกันมาก และในทางกลับกัน ค่าที่มากแสดงถึงองศาของมุมที่มาก หมายถึงความคล้ายคลึงกันน้อยนั่นเอง ซึ่งเป็นหลักการสำคัญที่ใช้ในเทคนิคการกรองข้อมูลตามเนื้อหา (Content-Based Filtering) และการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering)

2.4 การวัดประสิทธิภาพของแบบจำลอง (Evaluation Metrics)

2.4.1 ค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Error: MAE)

ค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Error: MAE) เป็นค่าที่ใช้ประเมินความแม่นยำของแบบจำลอง โดยคำนวณจากค่าความแตกต่างระหว่างค่าจริงและค่าที่คาดการณ์จากแบบจำลองแบบสัมบูรณ์ ($|Y_{pred} - Y_{true}|$) ของทุกจุดในแบบจำลอง แล้วนำค่าที่ได้มาหาค่าเฉลี่ยค่าผลลัพธ์ MAE จะเป็นบวกเสมอ เนื่องจากวัดจากค่าความคลาดเคลื่อนสัมบูรณ์ ค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ยที่ต่ำหมายถึงแบบจำลองสามารถคาดการณ์ได้ใกล้เคียงกับค่าจริงมาก (แบบจำลองมีความแม่นยำ) และหากค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ย เท่ากับ 0 แสดงว่าไม่มี ความคลาดเคลื่อนจากค่าจริง ค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ยมีสมการดังนี้

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_{pred,i} - y_{true,i}|$$

n = จำนวนตัวอย่างในชุดข้อมูล

$y_{pred,i}$ = ค่าที่คาดการณ์ลำดับที่ i

$y_{true,i}$ = ค่าความจริงลำดับที่ i

2.4.2 ค่าความคลาดเคลื่อนกำลังสองเฉลี่ย (Root Mean Square Error: RMSE)

เป็นตัวชี้วัดความคลาดเคลื่อน (error) ระหว่างค่าที่ทำนายได้จากแบบจำลองและค่าจริง โดย RMSE จะสะท้อนถึงค่าเบี่ยงเบนของผลการทำนายจากความเป็นจริง ยิ่งค่า RMSE ต่ำยิ่งแสดงว่าแบบจำลองสามารถทำนายค่าได้ใกล้เคียงกับความจริง และหากค่า RMSE เท่ากับ 0 หมายถึงไม่มีความคลาดเคลื่อนเลยในแบบจำลองนั้น

การใช้งาน RMSE ในระบบแนะนำ (Recommender System) โดยเฉพาะสำหรับการวัดผลการทำนายคะแนนที่ผู้ใช้น่าจะให้กับรายการ (Item) เช่น หากระบบแนะนำทำนายว่าผู้ใช้น่าจะให้เรตติ้ง 4 คะแนนกับภาพยนตร์เรื่องหนึ่ง แต่คะแนนจริงคือ 5 ค่าความคลาดเคลื่อนจะถูกนำมารวมกันและคำนวณค่า RMSE เพื่อวัดความแม่นยำในการทำนายโดยรวม ระบบแนะนำที่มีค่า RMSE ต่ำกว่าจึงบ่งบอกว่ามีความน่าเชื่อถือมากกว่าในการแนะนำรายการที่ผู้ใช้น่าจะชอบ

2.4.3 ค่าความแม่นยำที่ K (Precision@K)

เป็นตัวชี้วัดในระบบแนะนำ เพื่อประเมินความแม่นยำของผลลัพธ์ที่ถูกแนะนำใน K อันดับแรก โดยเฉพาะอย่างยิ่งในบริบทของการแนะนำรายการหรือผลิตภัณฑ์ของผู้ใช้ Precision@K จะช่วยตอบคำถามว่าในจำนวน K รายการที่แนะนำ มีรายการที่ผู้ใช้สนใจหรือเกี่ยวข้อง (Relevant Items) มากน้อยเพียงใด โดยคำนวณจากสัดส่วนของรายการที่เกี่ยวข้องกับ K อันดับแรก เทียบกับจำนวนรายการทั้งหมดใน K อันดับแรกนั้น ช่วงค่าของ Precision@K อยู่ระหว่าง 0-1 โดยค่าที่สูงแสดงถึงความแม่นยำที่ดี

โดยสมการของ Precision@K มีดังนี้

$$\text{Precision@K} = \frac{\text{จำนวนรายการที่เกี่ยวข้องใน K อันดับแรก}}{K}$$

จำนวนรายการที่เกี่ยวข้องใน K อันดับแรก = รายการที่เกี่ยวข้องหรือผู้ใช้สนใจ ซึ่งสามารถวัดได้จากข้อมูลเชิงพฤติกรรม (เช่น การคลิก, การซื้อ, การให้คะแนนสูง)

K = จำนวนอันดับต้นของรายการที่นำมาประเมิน (เช่น Top-5, Top-10)

2.4.3 ค่าความครบถ้วนที่ K (Recall@K)

เป็นตัวชี้วัดในระบบแนะนำ เพื่อประเมินความครบถ้วนของผลลัพธ์ที่ถูกแนะนำใน K อันดับแรก โดยจะช่วยตอบคำถามว่าในจำนวน K รายการที่แนะนำ มีกี่รายการที่เป็นรายการที่ผู้ใช้สนใจหรือเกี่ยวข้อง โดยคำนวณจากสัดส่วนของรายการที่เกี่ยวข้องใน K อันดับแรก เทียบกับ

รายการที่เกี่ยวข้องทั้งหมด ช่วงค่าของ Recall@K อยู่ระหว่าง 0-1 ยิ่งค่าใกล้ 1 ยิ่งแปลว่าระบบแนะนำสามารถดึงรายการที่เกี่ยวข้องได้ครบถ้วนมากขึ้น

สมการของ Recall@K มีดังนี้

$$\text{Recall@K} = \frac{\text{จำนวนรายการที่เกี่ยวข้องใน K อันดับแรก}}{\text{จำนวนรายการที่เกี่ยวข้องทั้งหมด}}$$

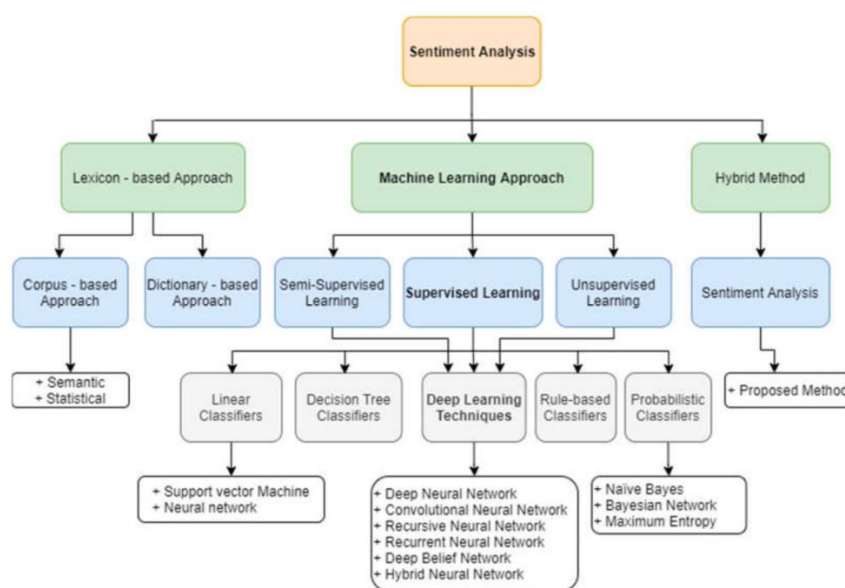
2.5 การวิเคราะห์ความรู้สึก (Sentiment Analysis)

การวิเคราะห์อารมณ์ (Sentiment Analysis) เป็นกระบวนการใช้คอมพิวเตอร์ในการวิเคราะห์และจำแนกความคิดเห็นของบุคคลที่แสดงออกผ่านรีวิว ความคิดเห็น หรือแบบสำรวจ โดยแบ่งเป็นอารมณ์เชิงบวก เชิงลบ หรือกลาง (Sahu et al., 2022)

การวิเคราะห์อารมณ์ (Sentiment Analysis) สามารถทำได้ใน 3 ระดับ ได้แก่ ระดับประโยค เอกสาร และแง่มุมหรือคุณลักษณะ โดยเป็นกระบวนการสกัดข้อมูลเพื่อระบุอารมณ์ของเอนทิตีอย่างอัตโนมัติ เพื่อวิเคราะห์ข้อความที่ผู้ใช้สร้างขึ้นว่าแสดงอารมณ์เชิงบวก เชิงลบ หรือกลาง

วิธีการวิเคราะห์อารมณ์แบ่งออกเป็น 3 แนวทางหลัก ได้แก่

1. เทคนิคที่อิงพจนานุกรม (Lexicon-Based Techniques): แบ่งเป็นแบบพจนานุกรม (Dictionary-Based) และแบบคอร์ปัส (Corpus-Based) ซึ่งเป็นวิธีเริ่มต้นในการจำแนกอารมณ์
2. เทคนิคที่อิงการเรียนรู้ของเครื่อง (Machine Learning-Based Techniques): รวมถึงการเรียนรู้แบบดั้งเดิมและการเรียนรู้เชิงลึก (Deep Learning)
3. วิธีผสมผสาน (Hybrid Approaches): ผสานเทคนิคการเรียนรู้ของเครื่องกับวิธีการอิงพจนานุกรม โดยพจนานุกรมอารมณ์มีบทบาทสำคัญในกระบวนการนี้



ภาพประกอบ 9 Taxonomy of sentiment analysis technique

ที่มา: https://www.researchgate.net/publication/354083599_An_Approach_to_Integrating_Sentiment_Analysis_into_Recommender_Systems

โดยงานวิจัยนี้ใช้การวิเคราะห์ข้อความจากผู้ใช้งานเพื่อจำแนกรีวิวออกเป็นอารมณ์เชิงบวกหรือลบจากรีวิวโดยใช้เทคนิคการเรียนรู้เชิงลึกเป็นหลัก

2.6 แบบจำลอง SieBERT

SieBERT ย่อมาจาก Sentiment in English RoBERTa ซึ่งเป็นหนึ่งในแบบจำลอง Transformer ที่พัฒนาต่อยอดมาจาก RoBERTa-large โดยใช้แนวคิดแบบ Transfer Learning โดยใช้การถ่ายโอนความเข้าใจภาษาที่ได้จากการฝึกฝนกับข้อมูลขนาดใหญ่ มาเพิ่มประสิทธิภาพสำหรับงานวิเคราะห์ความคิดเห็น (Sentiment Analysis) โดยเฉพาะ ถูกพัฒนาขึ้นโดย Jochen Hartmann, Mark Heitmann, Christian Siebert และ Christina Schamp

แบบจำลองนี้สามารถจำแนกความคิดเห็นเป็นเชิงบวกและเชิงลบได้อย่างมีประสิทธิภาพ โดยแบบจำลองนี้ผ่านการฝึกฝนและปรับแต่งโดยใช้ชุดข้อมูลความคิดเห็น (Sentiment-Labeled Datasets) จากแหล่งข้อมูลทั้งหมด 272 ชุด ซึ่งครอบคลุมข้อความที่มีการระบุอารมณ์มากกว่า 12 ล้านข้อความ และประเมินผลด้วยชุดข้อมูลที่หลากหลายถึง 15 ชุด เช่น รีวิวสินค้า ทวิต และข้อความทั่วไป เพื่อเพิ่มศักยภาพในการวิเคราะห์ความคิดเห็นจากแหล่งข้อมูลที่แตกต่างกัน โดยถูกออกแบบให้ใช้งานได้ง่ายเหมือนเครื่องมือสำเร็จรูป สามารถนำไปใช้งานได้ทั้งแบบไม่ต้อง

ปรับแต่งเพิ่มเติม (Pre-trained Model) และแบบปรับแต่งเพิ่มเติมสำหรับแอปพลิเคชันเฉพาะทาง (Fine-tune for Specific Tasks)

2.7 ค่าพารามิเตอร์ (Hyperparameter)

พารามิเตอร์เหล่านี้ใช้สำหรับการฝึกโมเดล โดยเน้นปรับแต่งการเรียนรู้ การบันทึกผล และการหยุดการฝึกในกรณีที่ไม่มีการพัฒนา

1. อัตราการเรียนรู้ (Learning Rate) สำหรับการปรับพารามิเตอร์ของโมเดล
2. จำนวนรอบ (Epochs) ที่โมเดลจะทำการฝึกกับข้อมูล
3. ขนาด Batch ที่ใช้สำหรับการฝึกบนแต่ละอุปกรณ์ (per_device_train_batch_size)
4. ขนาด Batch ที่ใช้สำหรับการประเมินผลบนแต่ละอุปกรณ์ (per_device_eval_batch_size)
5. จำนวนขั้นตอนสำหรับการ Warm-up Learning Rate ในช่วงเริ่มต้นการฝึก (warmup_steps)
6. ค่า Weight Decay สำหรับการปรับสมดุลระหว่างความแม่นยำของโมเดลและการป้องกันการ Overfitting
7. กลยุทธ์ในการบันทึกโมเดล เช่น บันทึกในแต่ละ Epoch (save_strategy="epoch")
8. กลยุทธ์สำหรับการประเมินผล เช่น ประเมินหลังจากจบแต่ละ Epoch (eval_strategy="epoch")
9. โหลดโมเดลที่ดีที่สุด (Best Model) เมื่อจบการฝึก (load_best_model_at_end= True)
10. ประเภทของ Learning Rate Scheduler ที่ใช้ เช่น Linear Scheduler (lr_scheduler_type)
11. ใช้ Callback เพื่อหยุดการฝึกก่อนกำหนด หากไม่มีการปรับปรุงค่าประสิทธิภาพ (เช่น Loss หรือ Accuracy) ในช่วง 5 Epoch ติดต่อกัน (early_stopping_patience)

2.8 การผสมผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำ

การผสมผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำ ด้วยวิธีการผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination)

การผสมผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำ (Recommender System) เป็นหนึ่งในแนวทางที่มีประสิทธิภาพในการเพิ่มความแม่นยำและความเหมาะสมของคำแนะนำ โดยเฉพาะในงานวิจัยที่มีจุดมุ่งหมายในการผสมการให้คะแนนสินค้า และผลจากการวิเคราะห์ความรู้สึกของลูกค้า

วิธีการผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) เป็นวิธีที่ได้รับความนิยมเนื่องจากความเรียบง่ายและประสิทธิภาพในการคำนวณ โดยมีข้อเสียคือต้องทำการทดลองหลายครั้งเพื่อหาค่าน้ำหนักที่เหมาะสม

วิธีนี้ใช้การรวมค่าคะแนนที่ได้จากแหล่งข้อมูลต่าง ๆ โดยคำนวณค่าเฉลี่ยแบบถ่วงน้ำหนัก (Weighted Average) เพื่อให้ได้ค่าที่สะท้อนความคิดเห็นของผู้ใช้งานได้อย่างแม่นยำที่สุด สูตรการคำนวณโดยทั่วไปคือ:

$$H = \alpha \cdot R + (1 - \alpha) \cdot S$$

โดยที่

H คือ ค่าคะแนนแบบผสมที่ใช้ในระบบแนะนำ

R คือ คะแนนเรตติ้งจากผู้ใช้งาน

S คือ คะแนนจากการวิเคราะห์ความรู้สึก

α คือ ค่าน้ำหนักความสำคัญที่กำหนดให้กับข้อมูลแต่ละประเภท โดยมีค่าอยู่ระหว่าง 0 ถึง 1

และอีกสองสมการที่ใช้ในงานวิจัยนี้เป็นการทดลองนำค่า probability มาใช้ด้วย เพื่อทดสอบประสิทธิภาพของระบบแนะนำ ได้แก่

$$H = \alpha \times R + (1 - \alpha) \times S \times \beta$$

และสมการ

$$H = \alpha \times R + (1 - \alpha) \times \beta$$

โดยที่

H คือ ค่าคะแนนแบบผสมที่ใช้ในระบบแนะนำ

R คือ คะแนนเรตติ้งจากผู้ใช้งาน

S คือ คะแนนจากการวิเคราะห์ความรู้สึก

α คือ ค่าน้ำหนักความสำคัญที่กำหนดให้กับข้อมูลแต่ละประเภท โดยมีค่าอยู่ระหว่าง 0 ถึง 1

β คือ ค่าความน่าจะเป็น (probability) ที่ได้จากการวิเคราะห์ความรู้สึก

2.9 การสร้างระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมแบบผสม

การสร้างระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมแบบผสม ด้วยวิธีการผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) เป็นการรวมผลลัพธ์ของระบบแนะนำการกรองรวมแบบอิงรายการ (Item-based) และแบบอิงผู้ใช้ (User-based) เพื่อสร้างคะแนนแนะนำที่สะท้อนความคิดเห็นหรือความชอบของผู้ใช้งานได้แม่นยำยิ่งขึ้น สูตรการคำนวณมีลักษณะดังนี้

$$P_{u,i} = \alpha \cdot CF_{u,i}^{user} + (1-\alpha) \cdot S_{u,i}$$

$P_{u,i}$: คะแนนแนะนำสำหรับผู้ใช้ u และสินค้า i

$CF_{u,i}^{user}$: คะแนนจากการกรองรวมแบบอิงผู้ใช้ (User-based)

$S_{u,i}$: คะแนนจากการกรองรวมแบบอิงรายการ (Item-based)

α : ค่าน้ำหนักที่กำหนดความสำคัญของข้อมูลแต่ละประเภท โดยมีค่าอยู่ระหว่าง 0 ถึง 1

2.10 งานวิจัยที่เกี่ยวข้อง (Related Research)

การศึกษางานวิจัยที่เกี่ยวข้อง ผู้วิจัยสนใจงานวิจัยที่ใช้แบบจำลองการวิเคราะห์ความรู้สึก เพื่อเพิ่มประสิทธิภาพให้กับระบบแนะนำ

2.10.1 บทความวิจัยเรื่อง An Approach to Integrating Sentiment Analysis into Recommender Systems (Dang et al., 2021)

บทความวิจัยนี้มุ่งเน้นการเพิ่มความน่าเชื่อถือให้กับระบบแนะนำแบบดั้งเดิม โดยนำเสนอการผสานกันระหว่างแบบจำลองการวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative filtering) โดยทำการทดลองบน 2 ชุดข้อมูล ได้แก่ Amazon Fine Foods Reviews และ Amazon Movie Reviews

สำหรับการวิเคราะห์ความรู้สึก (Sentiment Analysis) ผู้วิจัยได้เลือกใช้แบบจำลอง BERT (Bidirectional Encoder Representations from Transformers) ที่ผ่านการฝึกอบรมล่วงหน้า (pre-trained) ในการสร้างเวกเตอร์คุณลักษณะ (feature vector) จากข้อมูลรีวิว โดยเวกเตอร์คุณลักษณะที่ได้ถูกนำไปใช้ในการทดลองกับโครงข่ายประสาทเทียม (Neural Network) สองรูปแบบ คือ $BERT \rightarrow CNN \rightarrow LSTM$ (เรียกต่อไปว่า C-LSTM) และ $BERT \rightarrow LSTM \rightarrow CNN$ (เรียกต่อไปว่า L-CNN) เพื่อนำไปใช้ในการประเมินและเปรียบเทียบประสิทธิภาพของโมเดล

เพื่อเพิ่มประสิทธิภาพในการตรวจจับอารมณ์และรูปแบบข้อมูลในงานวิจัยนี้ ผู้วิจัยได้เลือกใช้ชั้น fully connected layer พร้อมกับ activation function แบบ ReLU (Rectified Linear Unit) เพื่อช่วยให้แบบจำลองสามารถเรียนรู้ข้อมูลเชิงลึกได้อย่างมีประสิทธิภาพมากขึ้น นอกจากนี้ยังได้มีการกำหนดป้ายกำกับ (label) ของข้อมูลในระดับอารมณ์ทั้งหมด 5 ระดับ ได้แก่ เชิงลบมาก (very negative), เชิงลบ (negative), เป็นกลาง (neutral), เชิงบวก (positive) และเชิงบวกมาก (very positive) ซึ่งสอดคล้องกับการจัดอันดับคะแนน (rating) ของระบบแนะนำ

การประเมินประสิทธิภาพของแบบจำลองวิเคราะห์ความรู้สึกใช้ตัวชี้วัดสามประเภท ได้แก่ ค่า Accuracy, ค่า AUC (Area Under the Curve) และค่า F1-score ซึ่งผลการทดลองแสดงให้เห็นว่า ทั้งสองแบบจำลองที่พัฒนาขึ้นสามารถให้ผลลัพธ์ที่น่าพึงพอใจ โดยมีค่า Accuracy มากกว่า 0.8, ค่า AUC สูงกว่า 0.84 และค่า F1-score เกิน 0.8 ในทุกชุดข้อมูล จากผลลัพธ์ดังกล่าว ผู้วิจัยจึงตัดสินใจเลือกนำแบบจำลองทั้งสองไปผสานกับระบบแนะนำ

ในระบบแนะนำ ผู้วิจัยและคณะเลือกใช้ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งาน (User-based collaborative filtering) และทำการทดสอบบน 3 อัลกอริทึมที่เป็นที่นิยม 3 ชนิด ได้แก่ Singular Value Decomposition (SVD), Non-Negative Matrix Factorization (NMF) และ SVD++ โดยใช้ระบบแนะนำที่ไม่ผสานการวิเคราะห์ความรู้สึกเป็นตัวเปรียบเทียบมาตรฐาน (baseline) และทดสอบการผสานระบบแนะนำร่วมกับแบบจำลองวิเคราะห์ความรู้สึกโดยกำหนดค่าน้ำหนัก (weight) ที่ 0.3, 0.5 และ 0.7 ตามลำดับ

ผลการทดสอบพบว่า การกำหนดค่าน้ำหนักที่ 0.3 ให้ผลลัพธ์ที่ดีที่สุดในทั้ง 3 อัลกอริทึม โดยชุดข้อมูล Amazon Movie Reviews ให้ผลลัพธ์ที่ดีกว่าชุดข้อมูล Amazon Fine Foods Reviews นอกจากนี้ยังพบว่า ผลการทดสอบระหว่างโมเดล L-CNN และ C-LSTM ไม่มีความแตกต่างอย่างมีนัยสำคัญในการใช้ร่วมกับระบบแนะนำแบบพึ่งพาผู้ใช้งาน

สำหรับการแนะนำแบบ Top-N Recommendation โดยใช้ MRR (Mean Reciprocal Rank), MAP (Mean Average Precision) และ NDCG (Normalized Discounted Cumulative Gain) เป็นตัวชี้วัด ผลลัพธ์แสดงให้เห็นว่าการใช้โมเดล C-LSTM ร่วมกับระบบแนะนำแบบพึ่งพาผู้ใช้งานและกำหนดค่าน้ำหนักที่ 0.7 ให้ค่าชี้วัดสูงสุด แต่การปรับปรุ้มนั้นมีเพียงเล็กน้อยและไม่เด่นชัดอย่างมีนัยสำคัญ

โดยสรุป การผสานแบบจำลองวิเคราะห์ความรู้สึกชนิด L-CNN และ C-LSTM กับระบบแนะนำแบบพึ่งพาผู้ใช้งานสามารถเพิ่มประสิทธิภาพของระบบได้จริง แต่พบว่าแบบจำลองวิเคราะห์ความรู้สึกทั้งสองแบบให้ผลลัพธ์ที่ใกล้เคียงกัน โดยค่าน้ำหนักที่เหมาะสมที่สุดคือ 0.3

อย่างไรก็ตาม การปรับปรุงประสิทธิภาพในแง่ของการแนะนำแบบ Top-N ยังไม่มีผลลัพธ์ที่โดดเด่น นอกจากนี้ ผลการทดสอบในชุดข้อมูล Amazon Movie Reviews ยังให้ประสิทธิภาพที่ดีกว่าชุดข้อมูล Amazon Fine Foods Reviews

2.10.2 บทความวิจัยเรื่อง Sentiment-aware Content Recommendation using LSTM-based Collaborative Filtering (Morzelona & Batra, 2022)

งานวิจัยนี้ได้นำเสนอแบบจำลองที่ผสมผสานเทคนิค Bi-LSTM, BERT และ Graph Convolutional Network (GCN) โดยมีวัตถุประสงค์เพื่อพัฒนาระบบแนะนำเนื้อหาที่สามารถวิเคราะห์และจับอารมณ์ของผู้ใช้ได้อย่างแม่นยำ ซึ่งเป็นประเด็นที่ระบบแนะนำทั่วไปมักประสบปัญหาในการตีความอารมณ์ที่ซ่อนอยู่ในข้อความที่ผู้ใช้สร้างขึ้น

ในขั้นตอนแรก ข้อมูลข้อความจะถูกแปลงเป็นเวกเตอร์ในมิติสูงด้วยเทคนิค word embeddings (แทนด้วย XWORD) จากนั้น XWORD จะถูกประมวลผลโดย Bi-LSTM และ BERT ในลักษณะคู่ขนาน Bi-LSTM มีหน้าที่จับลำดับการพึ่งพาและความหมายเชิงบริบท ในขณะที่ BERT จะจับความสัมพันธ์เชิงซับซ้อนระหว่างคำและประโยค ผลลัพธ์จากทั้งสองส่วนนี้จะถูกรวมกันในเลเยอร์การหลอมรวม (fusion layer) ซึ่งผลลัพธ์นี้จะถูกเรียกว่า E_fusion

ต่อมา E_fusion จะถูกป้อนเข้าสู่ Graph Convolutional Network (GCN) ซึ่งแบ่งออกการทำงานออกเป็นสองส่วนคือ Guser ที่ทำหน้าที่เรียนรู้โครงสร้างความสัมพันธ์ระหว่างผู้ใช้ และ Gcontent โดย Guser จะทำหน้าที่เรียนรู้โครงสร้างความสัมพันธ์ระหว่างผู้ใช้ และ Gcontent ที่เรียนรู้โครงสร้างความสัมพันธ์ของเนื้อหา ทั้งสองส่วนนี้จะประมวลผลแบบคู่ขนานเช่นกัน สุดท้าย embedding ที่ได้จาก GCN จะถูกนำไปใช้ในการสร้างคำแนะนำเนื้อหา

ผลการทดลองชี้ให้เห็นว่าแบบจำลองที่นำเสนอสามารถเพิ่มประสิทธิภาพได้อย่างมีนัยสำคัญ โดยค่า precision เพิ่มขึ้น 8.3%, accuracy เพิ่มขึ้น 8.5%, recall เพิ่มขึ้น 5.9%, speed เพิ่มขึ้น 4.5% และค่า AUC เพิ่มขึ้น 7.5% เมื่อเปรียบเทียบกับวิธีการอื่นที่มีอยู่

สรุปได้ว่า แบบจำลองที่นำเสนอมีความสามารถในการเพิ่มประสิทธิภาพของระบบแนะนำในทุกตัวชี้วัด ส่งผลให้ระบบสามารถตอบสนองต่อความต้องการและอารมณ์ของผู้ใช้ได้อย่างมีประสิทธิภาพมากขึ้น

2.10.3 บทความวิจัยเรื่อง Enhancing Collaborative Filtering-Based Recommender System Using Sentiment Analysis (Karabila, Darraz, El-Ansari, Alami, & El Mallahi, 2023)

งานวิจัยนี้นำเสนอระบบแนะนำเนื้อหาแบบใหม่ที่ผสมผสานการเรียนรู้แบบรวมกลุ่ม (Ensemble Learning) โดยการผนวกแบบผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกจากข้อความ รีวิวเข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) เพื่อ

ประยุกต์ใช้ในอุตสาหกรรม e-commerce โดยแบบจำลองที่ใช้สำหรับการวิเคราะห์ความรู้สึกคือ Bidirectional Long Short-Term Memory (Bi-LSTM) ซึ่งใช้ Global Vectors for Word Representation (GloVe) เป็น word embedding นอกจากนี้ ยังมีการปรับจูน hyperparameters และใช้เทคนิค dropout เพื่อลดปัญหาการ overfitting ส่งผลให้แบบจำลองวิเคราะห์ความรู้สึกมีค่า accuracy สูงถึง 0.93

การทดลองดำเนินการบนชุดข้อมูลจาก Amazon ที่มีขนาดแตกต่างกันสองชุด คือ Amazon Kindle Book (12,000 รายการ) และ Amazon Digital Music (64,000 รายการ) เพื่อทดสอบความสามารถของแบบจำลองในด้าน generalization และ robustness ในการประเมินประสิทธิภาพของระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวม (Collaborative Filtering) ผู้วิจัยและคณะได้เปรียบเทียบระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวม (Collaborative Filtering) สองรูปแบบ คือ การกรองรวมแบบอิงรายการ (Item-based) และแบบอิงผู้ใช้ (User-based) โดยใช้การวัดความคล้ายคลึงด้วย Cosine similarity

ผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึกถูกผสมผสานเข้ากับระบบแนะนำโดยการทดลองปรับค่าน้ำหนักตามสมการ $H = \alpha \times R + (1 - \alpha) \times S$ โดย H คือผลลัพธ์สุดท้าย, R คือคะแนนจากระบบแนะนำ, S คือคะแนนจากการวิเคราะห์ความรู้สึก, α คือค่าน้ำหนักที่ทดสอบโดยกำหนดค่าเป็น 0.3, 0.5 และ 0.7 ตามลำดับ โดยใช้ระบบแนะนำที่ไม่ผสมผสานการวิเคราะห์ความรู้สึกเป็นเป็นตัวเปรียบเทียบมาตรฐาน (baseline) การประเมินผลใช้ค่า Mean Absolute Error (MAE) และ Root Mean Square Error (RMSE) เป็นตัวชี้วัด

ผลการทดลองแสดงให้เห็นว่า ระบบแนะนำแบบอิงผู้ใช้ (User-based) ที่มีค่าน้ำหนัก 0.3 ให้ผลลัพธ์ที่ดีที่สุดในทั้งสองชุดข้อมูล โดยในชุด Amazon Kindle Book ค่า MAE ลดลงจาก 2.30 เป็น 1.12 และค่า RMSE ลดลงจาก 2.60 เป็น 1.28 สำหรับชุดข้อมูล Amazon Digital Music ค่า MAE ลดลงจาก 2.18 เป็น 1.15 และค่า RMSE ลดลงจาก 2.60 เป็น 1.28

สำหรับระบบแนะนำแบบอิงรายการ (Item-based) พบว่าค่าน้ำหนักที่เหมาะสมที่สุดคือ 0.7 โดยในชุดข้อมูล Amazon Kindle Book ค่า MAE ลดลงจาก 2.36 เป็น 1.18 และค่า RMSE ลดลงจาก 2.73 เป็น 1.35 สำหรับชุดข้อมูล Amazon Digital Music ค่า MAE ลดลงจาก 1.96 เป็น 1.32 และค่า RMSE ลดลงจาก 2.55 เป็น 1.46

สรุปได้ว่าการปรับค่าน้ำหนักให้เหมาะสม สามารถทำให้การผสมผสานแบบจำลองวิเคราะห์ความรู้สึกแบบ Bi-LSTM เข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวม (Collaborative Filtering) สามารถเพิ่มประสิทธิภาพของระบบแนะนำในชุดข้อมูล e-commerce ได้อย่างมีนัยสำคัญ

2.10.4 บทความวิจัยเรื่อง Recommendation system using Deep Learning-based sentiment analysis (Karabila, Darraz, El-Ansari, Alami, Lazaar, et al., 2023)

บทความวิจัยนี้มีวัตถุประสงค์เพื่อแก้ไขข้อจำกัดของระบบแนะนำแบบดั้งเดิมโดยการผสานแบบจำลองการวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำอย่างลงตัว สำหรับการวิเคราะห์ความรู้สึก ผู้วิจัยและคณะเลือกใช้แบบจำลอง Bidirectional Gated Recurrent Unit (Bi-GRU) ร่วมกับ Global Vectors for Word Representation (GloVe) เป็น word embedding ขณะที่ในส่วนของระบบแนะนำ ผู้วิจัยและคณะได้เลือกใช้เทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วม (Collaborative Filtering) ในสองรูปแบบ ได้แก่ การกรองข้อมูลรวมแบบอิงรายการ (Item-based) และการกรองข้อมูลรวมแบบอิงผู้ใช้ (User-based) โดยทั้งสองรูปแบบใช้การวัดความคล้ายคลึงด้วย Cosine similarity ดำเนินการทดสอบบนชุดข้อมูลรีวิวจาก Amazon Musical Instruments ซึ่งประกอบด้วยข้อมูลรีวิวจำนวน 10,262 รายการ

เหตุผลที่ผู้วิจัยเลือกใช้ Bi-GRU ในการวิเคราะห์ความรู้สึกเป็นเพราะสถาปัตยกรรมของ Bi-GRU มีความเรียบง่ายกว่า Bidirectional Long Short-Term Memory (Bi-LSTM) ซึ่งส่งผลให้การประมวลผลรวดเร็วขึ้น อีกทั้งยังช่วยลดปัญหา vanishing gradient ที่มักเกิดขึ้นในเครือข่ายประสาทแบบเวียนซ้ำ (Recurrent Neural Networks: RNNs) นอกจากนี้ เพื่อป้องกันปัญหาการเกิด overfitting ผู้วิจัยได้ใช้เทคนิคการปรับ regularization ด้วย dropout รวมถึงการหยุดการฝึกในเวลาที่เหมาะสม (Early stopping) และใช้เทคนิคการปรับพารามิเตอร์ด้วย Adaptive Moment Estimation (Adam) ซึ่งผลลัพธ์จากแบบจำลองการวิเคราะห์ความรู้สึกแสดงให้เห็นว่า ค่า Accuracy อยู่ที่ 0.89, ค่า AUC อยู่ที่ 0.76 และค่า F1-score อยู่ที่ 0.94

ในการบูรณาการแบบจำลองการวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำ ผู้วิจัยและคณะได้ให้ค่าน้ำหนักของผลลัพธ์จากการวิเคราะห์ความรู้สึกมากกว่าผลจากระบบแนะนำสองเท่า เนื่องจากมองว่าพฤติกรรมของผู้ใช้มักให้ความสำคัญกับบทวิจารณ์มากกว่าคะแนนความพึงพอใจ การทดสอบใช้ค่า Mean Absolute Error (MAE) และ Root Mean Square Error (RMSE) เป็นตัวชี้วัดและเปรียบเทียบผลลัพธ์กับระบบแนะนำที่ไม่มีการผสานแบบจำลองการวิเคราะห์ความรู้สึกเป็นตัวเปรียบเทียบมาตรฐาน (baseline)

ผลการทดสอบแสดงให้เห็นว่าในระบบแนะนำแบบอิงผู้ใช้ (User-based) ค่า MAE ลดลงจาก 2.50 เป็น 1.71 และค่า RMSE ลดลงจาก 2.80 เป็น 1.83 เมื่อเทียบกับ baseline ขณะที่ในระบบแนะนำแบบอิงรายการ (Item-based) ค่า MAE ลดลงจาก 2.06 เป็น 2.02 และค่า RMSE ลดลงจาก 2.60 เป็น 2.18

สรุปได้ว่าการบูรณาการแบบจำลองการวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำช่วยเพิ่มประสิทธิภาพของระบบแนะนำได้อย่างมีนัยสำคัญ โดยเฉพาะในระบบแนะนำแบบอิงผู้ใช้ (User-based) ซึ่งแสดงให้เห็นถึงการลดค่าความคลาดเคลื่อนได้อย่างชัดเจน

2.10.5 บทความวิจัยเรื่อง BERT-enhanced sentiment analysis for personalized e-commerce recommendations (Karabila, Darraz, EL-Ansari, Alami, & EL Mallahi, 2024)

งานวิจัยนี้มีจุดมุ่งหมายเพื่อแก้ไขข้อจำกัดของระบบแนะนำแบบดั้งเดิม โดยการผสานการวิเคราะห์อารมณ์ (sentiment analysis) เข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วม (collaborative filtering) เพื่อพัฒนาระบบแนะนำที่มีความแม่นยำและตรงตามความต้องการของผู้ใช้มากขึ้น โดยงานวิจัยนี้มุ่งเน้นที่การเพิ่มประสิทธิภาพของแบบจำลอง BERT ในการวิเคราะห์อารมณ์

สำหรับการพัฒนาแบบจำลองการวิเคราะห์อารมณ์ ผู้วิจัยเลือกใช้ Fine-tuned BERT ซึ่งเป็นแบบจำลองที่อยู่ในตระกูล Transformer โดยได้ทำการทดลองบนชุดข้อมูลรีวิว Amazon Musical Instruments จำนวน 10,262 รายการ โดยแบ่งข้อมูลเป็น training set 70% และ test set 30% การปรับแต่ง (Fine-tuning) ทำได้ 2 วิธี ได้แก่ การปรับแต่งด้วยตนเอง (manual) และการปรับแต่งอัตโนมัติ (automatic) ซึ่งในงานวิจัยนี้เลือกใช้วิธีการปรับแต่งด้วยตนเอง เนื่องจากเชื่อว่าการปรับด้วยตนเองช่วยให้สามารถเข้าใจโครงสร้างของโมเดลได้อย่างลึกซึ้งขึ้นและยังช่วยลดการใช้ทรัพยากรในการประมวลผล

กระบวนการ Fine-tuning ประกอบด้วยการใช้ BertTokenizer ซึ่งเป็น tokenizer เฉพาะสำหรับ BERT ในการทำ tokenization จากนั้นจึงเข้ารหัสข้อมูลด้วย pre-trained BERT เพื่อจับความหมายเชิงบริบทและความสัมพันธ์ระหว่างคำได้อย่างมีประสิทธิภาพ ผลลัพธ์ที่ได้คือ final hidden state ซึ่งแสดงถึงตัวแทนของแต่ละ token ในรีวิว จากนั้นจึงเพิ่ม Fully-Connected Layer เพื่อแปลงข้อมูลสำหรับการจำแนกประเภทแบบ Binary โดยใช้ฟังก์ชัน Sigmoid เพื่อทำนายผลลัพธ์ในช่วง 0 ถึง 1 โดยที่ 0 แสดงถึงอารมณ์เชิงลบ และ 1 แสดงถึงอารมณ์เชิงบวก ผลการทดลองแสดงว่า Fine-tuned BERT มีประสิทธิภาพสูงกว่าแบบจำลองอื่น ๆ โดยมีค่า AUC อยู่ที่ 0.87 เทียบกับ Bi-LSTM (0.69) และ Bi-GRU (0.76)

ในส่วนของพัฒนาระบบแนะนำ ผู้วิจัยและคณะเลือกใช้เทคนิค hybrid collaborative filtering ซึ่งเป็นการผสมผสานระหว่างการกรองข้อมูลร่วมแบบอิงผู้ใช้ (User-based) และการกรองร่วมข้อมูลแบบอิงรายการ (Item-based) โดยใช้แบบจำลอง Support Vector Regression (SVR) ในการรวมผลลัพธ์ ชุดข้อมูลสำหรับการทดสอบระบบแนะนำแบ่งเป็น 80% สำหรับการฝึกฝนและ 20% สำหรับการทดสอบ

ในการประเมินประสิทธิภาพของระบบแนะนำ ผู้วิจัยได้ทำการทดสอบระบบแนะนำแบบ User-based และ Item-based ร่วมกัน ผลการทดลองแสดงให้เห็นว่าระบบ Hybrid Collaborative Filtering ให้ผลลัพธ์ที่ดีที่สุด ผู้วิจัยจึงเลือกที่จะนำระบบนี้ไปผสานเข้ากับโมเดลการวิเคราะห์อารมณ์ในขั้นต่อไป

กระบวนการรวมผลลัพธ์จากการวิเคราะห์ความรู้สึกและระบบแนะนำใช้สมการ: $H = \alpha \times R + (1 - \alpha) \times S$ โดย H คือผลลัพธ์สุดท้าย, R คือคะแนนจากระบบแนะนำ, S คือคะแนนจากการวิเคราะห์ความรู้สึก และ α คือค่าน้ำหนักจากแบบจำลองวิเคราะห์ความรู้สึก ซึ่งผู้วิจัยและคณะได้ทำการทดลองค่าน้ำหนัก (α) ได้รับการทดสอบที่ค่า 0.3, 0.5, และ 0.7 ตามลำดับ

การประเมินผลประสิทธิภาพของระบบแนะนำนี้ใช้ตัวชี้วัด MAE และ RMSE ผลการทดสอบพบว่าระบบ Hybrid Collaborative Filtering ที่ผสานโมเดลการวิเคราะห์อารมณ์ร่วมด้วยที่ค่าน้ำหนัก $\alpha = 0.3$ ให้ผลลัพธ์ที่ดีที่สุด โดยค่า MAE ต่ำสุดอยู่ที่ 1.36 และค่า RMSE ต่ำสุดอยู่ที่ 1.43

ผลการวิจัยสรุปได้ว่า การผสานโมเดลการวิเคราะห์ความรู้สึกกับระบบแนะนำแบบ Hybrid Collaborative Filtering ช่วยเพิ่มประสิทธิภาพของระบบแนะนำได้อย่างมีนัยสำคัญ ทำให้คำแนะนำมีความน่าเชื่อถือและเฉพาะเจาะจงต่อความต้องการของผู้ใช้มากยิ่งขึ้น

ตาราง 1 แสดงผลการทดลองของงานวิจัยที่เกี่ยวข้อง

ลำดับ	ชื่อนักวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
1.	An Approach to Integrating Sentiment Analysis into Recommender Systems (Cach N. Dang, Maria N. Moreno-García and Fernando De la Prieta, 2021)	- Amazon Fine Food - Amazon Movie	แบบจำลองการวิเคราะห์ความรู้สึก (Sentiment Analysis) - Bert + LSTM + CNN (L-CNN) - Bert + CNN + LSTM (C-LSTM) ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวม (Collaborative Filtering)	- แบบจำลอง L-CNN ในชุดข้อมูล Amazon Fine Food ให้ผลการคำนวณ accuracy, F-score, AUC ที่ดีที่สุดคือ 80.04%, 80.24% และ 84.22% ตามลำดับ - แบบจำลอง L-CNN ในชุดข้อมูล Amazon Movie ให้ผลการคำนวณ accuracy, F-score, AUC ที่ดีที่สุดคือ 82.27%, 82.49% และ 86.07% ตามลำดับ	- อัลกอริทึม SVD ที่ทำการทดลองบนชุดข้อมูล Amazon Fine Food มีค่า MAE, RMSE และ NMAE ที่ดีที่สุดคือ 0.8365, 1.1338 และ 0.2091 ตามลำดับ โดยค่าที่ดีที่สุดมาจากการใช้แบบจำลอง L-CNN มารวมที่ค่านำหนัก 0.3 - อัลกอริทึม SVD ที่ทำการทดลองบนชุดข้อมูล Amazon Movie มีค่า MAE, RMSE และ NMAE ที่ดีที่สุดคือ 0.5943, 0.8732 และ 0.1486 ตามลำดับ โดยค่าที่ดีที่สุดมาจากการใช้แบบจำลอง L-CNN มารวมที่ค่านำหนัก 0.3 - อัลกอริทึม NMF ที่ทำการทดลองบนชุดข้อมูล Amazon Fine Food มีค่า MAE, RMSE และ NMAE ที่ดีที่สุดคือ 0.8762, 1.2102 (ที่ค่านำหนัก 0.5 โดยค่า RMSE ที่ค่านำหนัก 0.3 เท่ากับ 1.2103) และ 0.2191 ตามลำดับ จึง

ตาราง 1 (ต่อ)

ลำดับ	ชื่องานวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
				- แบบจำลอง C-LSTM ในชุดข้อมูล Amazon Movie ให้ผลการคำนวณ accuracy, F-score, AUC ที่ดีที่สุด คือ 82.27%, 82.46% และ 86.17% ตามลำดับ	สรุปว่าภาพรวมจากทุกตัวชี้วัด ค่าที่ดีที่สุดมาจากการใช้แบบจำลอง L-CNN มารวมที่ค่านำหนัก 0.3 - อัลกอริทึม NMF ที่ทำการทดลองบนชุดข้อมูล Amazon Movie มีค่า MAE, RMSE และ NMAE ที่ดีที่สุดคือ 0.5936, 0.9041 (ที่ค่านำหนัก 0.5 ส่วนที่ค่านำหนัก 0.3 RMSE มีค่าเท่ากับ 0.9112) และ 0.1484 ตามลำดับ โดยค่าที่ดีที่สุดมาจากการใช้แบบจำลอง L-CNN มารวมที่ค่านำหนัก 0.3 - อัลกอริทึม SVD++ ที่ทำการทดลองบนชุดข้อมูล Amazon Fine Food มีค่า MAE, RMSE และ NMAE ที่ดีที่สุดคือ 0.8263, 1.1292 และ 0.2066 ตามลำดับ โดยค่าที่ดีที่สุดมาจากการใช้แบบจำลอง L-CNN มารวมที่ค่านำหนัก 0.3 - อัลกอริทึม SVD++ ที่ทำการทดลองบนชุดข้อมูล Amazon Movie มีค่า MAE, RMSE และ NMAE ที่ดีที่สุดคือ 0.5770, 0.0.8577 และ

ตาราง 1 (ต่อ)

ลำดับ	ชื่องานวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
				0.1443ตามลำดับ โดยค่าที่ดีที่สุดมาจากการใช้แบบจำลอง L-CNN มารวมที่ค่าน้ำหนัก 0.3	
				ผู้วิจัยสรุปว่าการใช้วิธีผสมผสานแบบจำลองวิเคราะห์ความรู้สึกเข้ากับระบบแนะนำให้ผลดีกว่าในทุกค่าน้ำหนักเมื่อเทียบกับการใช้ระบบแนะนำเพียงอย่างเดียว โดยให้ผลลัพธ์ใกล้เคียงกันทั้งในแบบจำลอง C-LSTM และ L-CNN คือค่าน้ำหนักที่ดีที่สุดคือ 0.3	
					ในการแนะนำสินค้าแบบ top-N ในชุดข้อมูล Amazon Fine Food ที่ใช้แบบจำลอง C-LSTM ที่ค่าน้ำหนัก 0.7 ช่วยเพิ่มค่า MAP ได้ 0.58 % สำหรับอัลกอริทึม SVD และ 0.36% สำหรับอัลกอริทึม SVD++ ซึ่งถือว่าไม่มีนัยสำคัญ

ตาราง 1 (ต่อ)

ลำดับ	ชื่อนักวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
2.	Sentiment-aware Content Recommendation using LSTM-based Collaborative Filtering (Prof. Romi Morzelona and Mrs. Sweta Batra, 2022)	มี 3 ชุดข้อมูล แต่ผู้วิจัยไม่ได้ระบุไว้	- Bi-LSTM - BERT - Graph Convolutional Network (GCN)	-	ค่า precision เพิ่มขึ้น 8.3%, accuracy เพิ่มขึ้น 8.5%, recall เพิ่มขึ้น 5.9%, speed เพิ่มขึ้น 4.5% และค่า AUC เพิ่มขึ้น 7.5% เมื่อเปรียบเทียบกับวิธีการที่มีอยู่
3.	Enhancing Collaborative Filtering-Based Recommender System Using Sentiment Analysis (Ikram Karabila, Nossayba Darraz, Anas El-Ansari, Nabil Alami and Mostafa El Mallahi, 2023)	- Amazon Kindle book (12,000 แถว) - Amazon digital music (64,000 แถว)	- Bi-LSTM - Collaborative Filtering (User-based) - Collaborative Filtering (Item-based)	- จากการทดลองในชุดข้อมูลรีวิว Amazon Kindle book ได้ผลลัพธ์ดังนี้ ค่า Accuracy ที่ 0.93 ค่า F1-Score ที่ 0.94 และ ค่า AUC ที่ 0.92 2. Amazon digital music ซึ่งเป็นชุดข้อมูลรีวิวที่มีขนาดใหญ่กว่าได้ผลลัพธ์ดังนี้ ค่า Accuracy ที่ 0.94 และ F1-Score ที่ 0.96 อย่างไรก็ตาม ค่าดัชนี AUC ลดลงมาอยู่ที่ 0.78	- ผลการวัดประสิทธิภาพของระบบแนะนำแบบ User-based Collaborative Filtering บนชุดข้อมูล Amazon Kindle Book พบว่าการผสานแบบจำลองวิเคราะห์ความรู้สึกที่คำนวณจาก 0.7 ให้ผลลัพธ์ที่ดีที่สุด โดยมีค่า MAE เท่ากับ 1.12 และค่า RMSE เท่ากับ 1.30 ซึ่งดีกว่าการไม่ใช้แบบจำลองวิเคราะห์ความรู้สึกซึ่งมีค่า MAE เท่ากับ 2.30 และค่า RMSE เท่ากับ 2.63 - ผลการวัดประสิทธิภาพของระบบแนะนำแบบ User-based Collaborative Filtering บนชุดข้อมูล Amazon Digital Music พบว่าการผสาน

ลำดับ	ชื่องานวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
				แบบจำลองวิเคราะห์ความรู้สึกที่ค่าน้ำหนัก 0.7 ให้ผลลัพธ์ที่ดีที่สุด โดยมีค่า MAE เท่ากับ 1.15 และค่า RMSE เท่ากับ 1.28 ซึ่งดีกว่าการไม่ใช้แบบจำลองวิเคราะห์ความรู้สึกที่มีค่า MAE เท่ากับ 2.18 และค่า RMSE เท่ากับ 2.66 - ผลการวัดประสิทธิภาพของระบบแนะนำแบบ item-based collaborative filtering (CF) บนชุดข้อมูล Amazon Digital Music พบว่า การผสมระบบแนะนำเข้ากับแบบจำลองการวิเคราะห์ความรู้สึก (มีน้ำหนัก 7) มีค่า Mean Absolute Error (MAE) ที่ 1.18 และ Root Mean Squared Error (RMSE) ที่ 1.35 ซึ่งดีกว่าการใช้ระบบแนะนำแบบเดิมโดยไม่ใช้การวิเคราะห์ความรู้สึก (ที่มีค่า MAE 2.36 และ RMSE 2.73) - ผลการวัดประสิทธิภาพของระบบแนะนำแบบ item-based collaborative filtering (CF) บนชุดข้อมูล Amazon Digital Music พบว่า การรวม	

ตาราง 1 (ต่อ)

ลำดับ	ชื่องานวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
4.	Recommendation system using Deep Learning-based sentiment analysis (Ikram Karabila, Nossayba Darraz, Anas El-Ansari, Nabil Alami, Mohamed LAZAAR and Mostafa El Mallahi, 2023)	- Amazon Musical Instruments reviews dataset	- BiGRU -Collaborative Filtering (User-based) -Collaborative Filtering (Item-based) โดยได้ทำการปรับค่าน้ำหนักจากแบบจำลองวิเคราะห์ความรูสึกเป็น 2 เท่าของระบบแนะนำ	- Accuracy 0.89 - AUC 0.76 - F1-Score 0.94	ระบบแนะนำกับแบบจำลองการวิเคราะห์ความรูสึกที่ค่าน้ำหนัก 0.3 ให้ผลลัพธ์ที่ดีที่สุด โดยมีค่า Mean Absolute Error (MAE) เท่ากับ 1.32 และ Root Mean Squared Error (RMSE) เท่ากับ 1.46 เมื่อเปรียบเทียบกับการใช้ระบบแนะนำโดยไม่มีการวิเคราะห์ความรูสึก ซึ่งมีค่า MAE เท่ากับ 21.96 และ RMSE เท่ากับ 2.55
					- ระบบแนะนำแบบ user-based : กรณีที่ไม่ได้รวมข้อมูลจากแบบจำลองวิเคราะห์ความรูสึก ค่า MAE เท่ากับ 2.50 และค่า RMSE เท่ากับ 2.80 - ระบบแนะนำแบบ user-based : กรณีที่รวมข้อมูลจากแบบจำลองวิเคราะห์ความรูสึก ค่า MAE เท่ากับ 1.71 และค่า RMSE เท่ากับ 1.83 - ระบบแนะนำแบบ item-based : กรณีที่ไม่ได้รวมข้อมูลจากแบบจำลองวิเคราะห์ความรูสึก ค่า MAE เท่ากับ 2.06 และค่า RMSE เท่ากับ 2.60

ตาราง 1 (ต่อ)

ลำดับ	ชื่องานวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
5.	BERT-enhanced sentiment analysis for personalized e-commerce recommendations (Ikram Karabila, Nossayba Darraz, Anas El-Ansari, Nabil Alami and Mostafa El Mallahi, 2023)	- Amazon Musical Instruments reviews dataset	- Fine-tuned BERT - Collaborative Filtering (User-based) - Collaborative Filtering (Item-based) - Hybrid Collaborative Filtering (User-based, Item-based)	- Accuracy 0.91 - AUC 0.87 - F1-Score 0.95 - MCC 0.50	- ระบบแนะนำแบบ user-based : กรณีสหรวมข้อมูลด้านจากแบบจำลองวิเคราะห์ความรู้สึก ค่า MAE เท่ากับ 2.00 และค่า RMSE เท่ากับ 2.18 - กรณีสหรวมข้อมูลจากแบบจำลองวิเคราะห์ความรู้สึก - User-based CF: ค่า RMSE เท่ากับ 2.80 และค่า MAE เท่ากับ 2.50 - Item-based CF: ค่า RMSE เท่ากับ 2.60 และค่า MAE เท่ากับ 2.06 - Hybrid CF: ค่า RMSE เท่ากับ 2.00 และค่า MAE เท่ากับ 1.58 - กรณีสหรวมข้อมูลแบบ Hybrid CF มากรวมข้อมูลกับแบบจำลองวิเคราะห์ความรู้สึก - ไม่รวมข้อมูล SA: ค่า RMSE เท่ากับ 2.00 และค่า MAE เท่ากับ 1.58

ตาราง 1 (ต่อ)

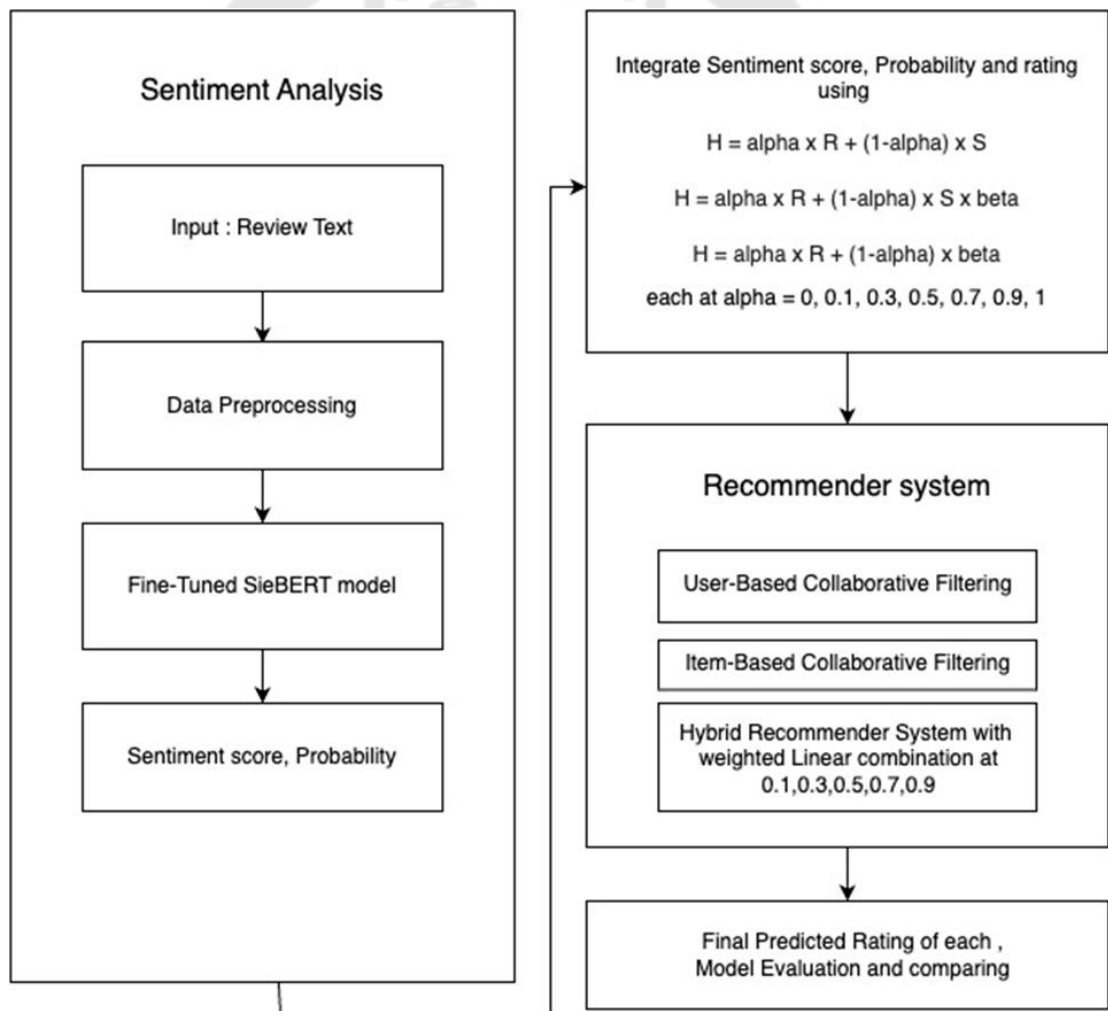
ลำดับ	ชื่องานวิจัย	ชุดข้อมูล	แบบจำลอง	ผลลัพธ์ Sentiment Analysis Model (SA)	ผลลัพธ์ Collaborative filtering + SA
					- รวมข้อมูล SA ที่คำนวณเท่ากับ 3 ค่า RMSE ลดลงเหลือ 1.43 และค่า MAE ลดลงเหลือ 1.36 - รวมข้อมูล SA ที่คำนวณเท่ากับ 5: ค่า RMSE ลดลงเหลือ 1.44 และค่า MAE ลดลงเหลือ 1.37 - รวมข้อมูล SA ที่คำนวณเท่ากับ 7: ค่า RMSE เท่ากับ 1.47 และค่า MAE เท่ากับ 1.39

บทที่ 3

การดำเนินงานวิจัย

ในการดำเนินงานวิจัย ผู้วิจัยได้ดำเนินการดังภาพประกอบ 10 โดยมีขั้นตอนดังต่อไปนี้

1. การประมวลผลข้อมูลเบื้องต้น
2. การสร้าง ปรับค่าพารามิเตอร์ของแบบจำลอง และทำ Sentiment Analysis
3. การสร้างฟีเจอร์จาก Sentiment Analysis
4. การนำชุดข้อมูลที่ได้มาเข้ากับระบบแนะนำ
5. การอภิปรายและสรุปผลการดำเนินงาน



ภาพประกอบ 10 แผนภาพแสดงกระบวนการดำเนินงานของงานวิจัยนี้

3.1 การประมวลผลข้อมูลเบื้องต้น

การวิจัยครั้งนี้ทางผู้วิจัยได้นำข้อมูล Amazon Kindle Review จากแหล่งข้อมูลสาธารณะ โดยใช้ไฟล์ข้อมูล Kindle_Store.csv ซึ่งประกอบด้วยข้อมูลจำนวน 17,383,985 แถว 10 คอลัมน์ ดังภาพประกอบ 11

คอลัมน์	ประเภทข้อมูล	คำอธิบาย
rating	float	คะแนนของสินค้า (ตั้งแต่ 1.0 ถึง 5.0)
title	str	หัวข้อของรีวิวผู้ใช้งาน
text	str	เนื้อหาของรีวิวที่ผู้ใช้งานเขียน
images	list	ภาพที่ผู้ใช้งานโพสต์หลังจากได้รับสินค้า แต่ละภาพมีขนาดต่างกัน (เล็ก, กลาง, ใหญ่) โดยใช้ small_image_url, medium_image_url และ large_image_url ตามลำดับ
asin	str	รหัสของสินค้า
parent_asin	str	รหัสพ่อแม่ของสินค้า หมายถึง: สินค้าที่มีสี, สไตล์, หรือขนาดต่างกันมักจะอยู่ภายใต้รหัสพ่อแม่เดียวกัน (asin ในชุดข้อมูล Amazon ก่อนหน้านี้คือรหัสพ่อแม่) ใช้รหัสพ่อแม่เพื่อค้นหาหมวดหมู่ของสินค้า
user_id	str	รหัสของผู้รีวิว
timestamp	int	เวลารีวิว (รูปแบบ unix time)
verified_purchase	bool	การยืนยันการซื้อสินค้าของผู้ใช้งาน
helpful_vote	int	คะแนนโหวตว่ามีประโยชน์ของรีวิว

ภาพประกอบ 11 โครงสร้างข้อมูล Kindle_Store.csv

3.1.1 การเตรียมข้อมูลให้พร้อมใช้งาน

ทำการกรองข้อมูลเฉพาะคอลัมน์ที่จำเป็นเพื่อนำมาวิเคราะห์ ได้แก่ คอลัมน์ ผู้ใช้งาน (user_id) รหัสสินค้าหรือรายการ (parent_asin) หัวข้อของรีวิวผู้ใช้งาน (title) เนื้อหารีวิวที่ผู้ใช้งานเขียน (text) และการให้คะแนนจากผู้ใช้งานจาก 1-5 คะแนน (rating) จากนั้นทำการกรองข้อมูลซ้ำซ้อน ข้อมูลที่หนึ่งผู้ใช้งานรีวิวรายการเดียวกันมากกว่า 1 ครั้ง โดยเก็บเฉพาะครั้งล่าสุดที่รีวิว รวมถึงกรองข้อมูลที่ผู้ใช้คนเดียวรีวิวน้อยกว่า 20 รีวิวออก ก่อนจะรวมคอลัมน์ หัวข้อของรีวิวผู้ใช้งาน (title) และเนื้อหารีวิวที่ผู้ใช้งานเขียน (text) เป็นคอลัมน์เดียว หลังจากนั้นได้ทำ label ของข้อมูลโดยให้คะแนนเรตติ้ง 1-3 เป็นค่าลบ (NEGATIVE) และ คะแนนเรตติ้ง 4-5 เป็นค่าบวก (POSITIVE) เนื่องจากพบว่ารีวิวที่มีค่าเท่ากับ 3 บ่งบอกถึงความไม่ประทับใจในหนังสือมากนัก อาทิ “There

are NO recipe pics. Such a disappointment. Borrow for free if you can... no really innovative recipes. Could be a gift for 10 yr old learning to cook?? Honestly 3 stars is generous.” หรือ “I liked it but probably not enough to continue the series. For me there wasn't enough action, not enough romance...just not enough to capture my attention. Good character back stories but nothing to make me fall in love with them. I liked the other characters, even if there were so many to keep track of. Love the author but not this book.” หรือ “Too predictable. Not one of the best by far Too slow moving story. Dragging.” เป็นต้น

จากนั้นได้เปลี่ยนชื่อคอลัมน์เป็น ผู้ใช้งาน (USER) รหัสสินค้าหรือรายการ (ITEM) หัวข้อของรีวิวผู้ใช้งาน (REVIEW) การให้คะแนนจากผู้ใช้งานจาก 1-5 คะแนน (RATING) และ คอลัมน์ป้ายกำกับเป็น TRUE_SENTIMENT เมื่อทำการกรอกร่องข้อมูลเฉพาะคอลัมน์ที่จำเป็นเรียบร้อยแล้ว จะได้ผลลัพธ์ดังภาพประกอบ 12

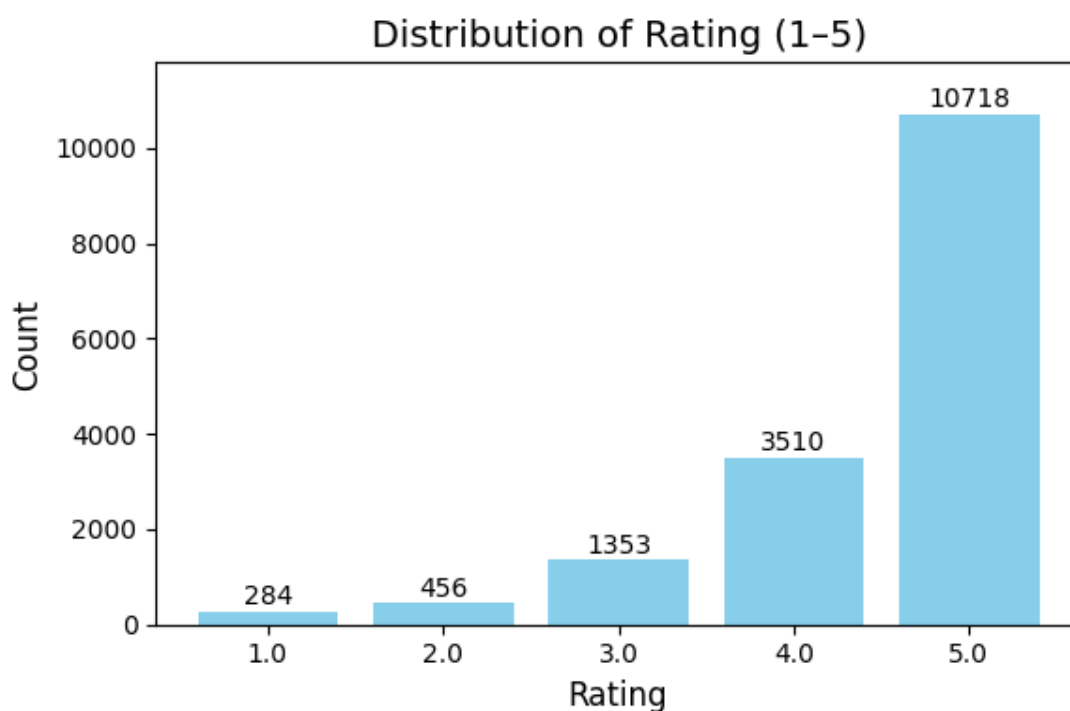
	USER	ITEM	TITLE	REVIEW	RATING	TRUE_SENTIMENT
0	AHGTHCERTEZUXNBLJ5SWHK2CDLXA	B073DFP8VC	Alright book	The book Fade was not my favorite but was a go...	3.0	NEGATIVE
1	AHFY2QSS6PK5MHSYZF6TXUYNPLQ	B07QVH25KX	Hats off to Fern Michaels for all her great ac...	I have been a fan of this author for many year...	5.0	POSITIVE
2	AFWHJ6O3PV4JC7PVOJH6CPULO2KQ	B08M993CNC	Add to legend f Arthur	Good twist on the ledgen of King Arthur l...	5.0	POSITIVE
3	AFWHJ6O3PV4JC7PVOJH6CPULO2KQ	B004NNVWL6	MEMORIES	Brought back a lot of memories. Reading sci-f...	4.0	POSITIVE
4	AFWHJ6O3PV4JC7PVOJH6CPULO2KQ	B094CHNKCK	What a woman!	Enjoyed the story! Marveled at the characters ...	5.0	POSITIVE
...	...	...	...	...	...	...
17383980	AFB25U4WVVKZMJJZCKHPFXSKL7SQ	B0118DAWR8	Four Stars	Not bad, could have more sequels	4.0	POSITIVE
17383981	AH4PJ73QN75AJM5VSCT53AOADCGA	B004PYDNFG	SPECTACULAR	This is one of my most cherished books of all ...	5.0	POSITIVE
17383982	AGTABH5ZMAB3DTPJMJYVSPPMKMA	B00D7Z4GQY	Love this book	I read this for a book club. Absolutely loved ...	5.0	POSITIVE
17383983	AFKC3IC2B3LDMJMDYY4UHDQPYAVA	B01MA5I6QU	I loved sharing this with my daughter	I loved sharing this book with my daughter who...	5.0	POSITIVE

ภาพประกอบ 12 ผลลัพธ์เมื่อเตรียมข้อมูลเรียบร้อยแล้ว

3.1.2 การแบ่งชุดข้อมูล

ชุดข้อมูลได้ถูกแบ่งออกเป็น 4 ชุด โดยแบ่งตามผู้ใช้งานเป็นหลัก โดยผู้ใช้งานหนึ่งคนจะปรากฏอยู่ใน 1 ชุดข้อมูลเท่านั้น ดังนี้

1. ชุดข้อมูลการฝึก (Training set) - 70% หรือ 1,400 ผู้ใช้งานที่ไม่ซ้ำกัน จำนวน 2,071 รีวิว โดยเนื่องจากเป็นชุดข้อมูลสำหรับฝึกแบบจำลอง จึงถูกแบ่งอย่างสมดุล โดยแบ่งออกเป็น 1,186 NEGATIVE (56.5%) และ 912 POSITIVE (43.5%)
2. ชุดข้อมูลตรวจสอบ (Validation set) - 15% หรือ 300 ผู้ใช้งานที่ไม่ซ้ำกัน จำนวน 1,021 รีวิว
3. ชุดข้อมูลทดสอบ (Test set) - 15% หรือ 300 ผู้ใช้งานที่ไม่ซ้ำกัน จำนวน 986 รีวิว
4. ชุดข้อมูลสำหรับพยากรณ์ (Predict set) - 5,000 ผู้ใช้งานที่ไม่ซ้ำกัน ใช้สำหรับพยากรณ์ความรู้สึกและรวมเข้ากับคะแนนรีวิวเพื่อนำไปใช้กับระบบแนะนำ จำนวน 16,321 รีวิว แบ่งออกเป็น 2,093 NEGATIVE (12.8%) และ 14,228 POSITIVE (87.2%) โดยการกระจายตัวของข้อมูลเป็นดังภาพประกอบ 13



ภาพประกอบ 13 การกระจายตัวของข้อมูลในชุด Predict set

โดยชุด Predict set จะถูกแบ่งออกเป็น 2 แบบ ได้แก่ ชุดข้อมูลที่มีเรตติ้ง 1-5 และชุดข้อมูลที่แบ่งเรตติ้งออกเป็น 0 และ 1 โดยให้เรตติ้งตั้งแต่ 1-3 เป็น 0 และ 4-5 เป็น 1

3.1.3 การประมวลผลล่วงหน้า (preprocessing) ของข้อมูลรีวิว

ในการประมวลผลข้อมูลล่วงหน้า (preprocessing) ได้ดำเนินการกับข้อมูลทั้ง 4 ชุด เพื่อเตรียมข้อมูลให้อยู่ในรูปแบบที่เหมาะสมสำหรับการวิเคราะห์และการป้อนเข้าสู่โมเดล กระบวนการเริ่มต้นด้วยการลบแท็ก HTML ที่ไม่จำเป็นออกจากข้อความ (Remove HTML Tag) ซึ่งช่วยลดความซับซ้อนของข้อมูล จากนั้นได้ทำการขยายคำย่อให้เป็นคำเต็ม (Expanding Contraction) เพื่อรักษาความหมายของข้อความให้สมบูรณ์และชัดเจน

นอกจากนี้ ยังได้แก้ไขการสะกดคำที่ผิดพลาด (Spelling Correction) และแปลงข้อความทั้งหมดให้เป็นตัวพิมพ์เล็ก (Lowercase) เพื่อให้การประมวลผลมีความสม่ำเสมอ รวมถึงจัดการกับลิงก์ URL, การกล่าวถึงผู้ใช้งาน (@Mentions) และแฮชแท็ก (Handling URLs, Mentions, and Hashtags) ซึ่งเป็นส่วนที่อาจไม่เกี่ยวข้องกับเนื้อหาหลักของข้อความ

ในกระบวนการถัดมา ได้ลบตัวเลขทั้งหมดออกจากข้อความ (Remove all numbers from the text) และลบเครื่องหมายวรรคตอนที่ไม่มีความสำคัญในเชิงวิเคราะห์ (Remove Punctuation) พร้อมทั้งจัดการกับช่องว่างที่เกินความจำเป็น (Remove Extra Spaces) เพื่อปรับโครงสร้างข้อความให้กระชับและเหมาะสมสำหรับการใช้งาน

เพื่อปรับปรุงคุณภาพของข้อมูลเพิ่มเติม ได้ลบคำที่ไม่มีความสำคัญ เช่น คำที่ซ้ำบ่อย ในประโยคแต่ไม่ได้มีความหมายเชิงวิเคราะห์ (Remove Stop Words) และทำการแปลงคำให้อยู่ในรูปแบบรากศัพท์ (Lemmatization) เพื่อให้แต่ละคำอยู่ในรูปแบบที่เป็นมาตรฐาน กระบวนการทั้งหมดนี้ช่วยให้ข้อมูลมีความสะอาดและพร้อมสำหรับการนำไปใช้งานในขั้นตอนถัดไป

3.2 การเทรน ปรับแต่ง และประเมินผลแบบจำลอง SieBERT สำหรับการวิเคราะห์ความรู้สึก

ในการปรับแต่งและประเมินผลโมเดล SieBERT สำหรับการวิเคราะห์ความรู้สึก เริ่มต้นด้วยการเตรียมชุดข้อมูล โดยแปลงข้อมูลในรูปแบบ DataFrame เช่น train_df, val_df และ test_df ให้อยู่ในรูปแบบ Hugging Face Dataset เพื่อให้รองรับการประมวลผลได้สะดวกยิ่งขึ้นผ่านฟังก์ชัน Dataset.from_pandas() หลังจากนั้น ได้สร้างฟังก์ชัน tokenize_function เพื่อแปลงข้อความในคอลัมน์ REVIEW ให้เป็น Tokenized Input โดยใช้ Tokenizer ของ SieBERT ซึ่งฟังก์ชันนี้รวมถึงการจัดการข้อความยาวเกินขนาดที่กำหนด (truncation) และการเพิ่ม Padding เพื่อให้ข้อมูลมีความยาวเท่ากัน

ฟังก์ชัน preprocess_function ถูกพัฒนาขึ้นเพื่อเตรียมข้อมูลเพิ่มเติม โดยทำการเพิ่มป้ายกำกับ (labels) ที่แปลงจากคะแนนรีวิว (rating) ให้เป็นค่าเชิงบวกหรือเชิงลบ (0 หรือ 1) ตามเงื่อนไขที่กำหนด พร้อมทั้งนำข้อมูลมาทำ Tokenization ชุดข้อมูลการฝึก (training set) และชุด

ตรวจสอบ (validation set) ถูกนำมาประมวลผลผ่านฟังก์ชันดังกล่าวด้วยวิธี `.map()` เพื่อให้พร้อมสำหรับการฝึกโมเดล ในส่วนของฮาร์ดแวร์ โมเดลถูกย้ายไปยังอุปกรณ์ที่เหมาะสม เช่น GPU (cuda) หรือ CPU ขึ้นอยู่กับความพร้อมของระบบ

เพื่อประเมินประสิทธิภาพของโมเดล ได้กำหนดฟังก์ชัน `compute_metrics` สำหรับคำนวณตัวชี้วัดที่สำคัญ ได้แก่ ความแม่นยำ (accuracy), ค่า F1-score และค่าพื้นที่ใต้กราฟ ROC (AUC) ซึ่งช่วยในการวัดประสิทธิภาพของโมเดลในการจำแนกข้อมูลอย่างละเอียด นอกจากนี้ ยังได้ตั้งค่าพารามิเตอร์การฝึกโมเดลผ่าน `TrainingArguments` โดยกำหนดอัตราการเรียนรู้ (learning rate) เป็น $5e-6$ จำนวนรอบการฝึก (epochs) เป็น 50 และขนาด Batch สำหรับการฝึกและการตรวจสอบเป็น 16 พร้อมเปิดใช้งาน Early Stopping หากไม่มีการปรับปรุงประสิทธิภาพใน 5 รอบการฝึก รวมถึงตั้งค่าให้บันทึกโมเดลที่ดีที่สุดและทำการประเมินผลในทุก epoch

กระบวนการฝึกและประเมินผลโมเดลดำเนินการผ่าน Trainer ของ Hugging Face โดยตั้งค่าโมเดล ชุดข้อมูล ตัวชี้วัด และพารามิเตอร์ต่าง ๆ เมื่อเริ่มต้นการฝึกโมเดลด้วยคำสั่ง `trainer.train()` โมเดลจะถูกปรับแต่งตามข้อมูลการฝึก และประเมินผลบนชุดข้อมูลตรวจสอบ ผลลัพธ์ที่ได้อ้างอิงถึงค่าความแม่นยำ (accuracy), ค่า AUC และ Classification Report ซึ่งบ่งชี้ถึงประสิทธิภาพของโมเดล

สุดท้าย โมเดลที่ปรับแต่งเสร็จแล้วและตัว Tokenizer ถูกบันทึกไว้ในตำแหน่งที่กำหนด พร้อมทั้งบันทึกผลลัพธ์จากชุดข้อมูลตรวจสอบในรูปแบบไฟล์ CSV โดยเพิ่มคอลัมน์ `TRUE_SENTIMENT` (ค่าจริง) และ `PREDICTED_SENTIMENT` (ค่าที่โมเดลคาดการณ์) กระบวนการเหล่านี้ช่วยให้โมเดล SieBERT สามารถวิเคราะห์ความรู้สึกจากข้อความได้อย่างมีประสิทธิภาพ พร้อมทั้งสร้างผลลัพธ์ที่สามารถนำไปใช้หรือวิเคราะห์ต่อได้ในขั้นตอนถัดไป

3.3 การสร้างพีเจอร์จาก Sentiment Analysis

หลังจากที่แบบจำลอง SieBERT ได้รับการฝึกและปรับแต่งจนได้ผลลัพธ์ที่น่าพึงพอใจในขั้นตอนก่อนหน้านี้ ได้มีการนำแบบจำลองนี้มาใช้กับชุดข้อมูลสำหรับพยากรณ์ (Predict set) กระบวนการนี้มีเป้าหมายเพื่อให้แบบจำลองคาดการณ์ค่าความรู้สึก (sentiment) จากข้อความรีวิวในรูปแบบค่าทางตัวเลข โดยค่าที่ได้จะประกอบด้วยป้ายกำกับ (positive หรือ negative) และค่าความน่าจะเป็นที่แสดงถึงความมั่นใจในผลลัพธ์ของแบบจำลอง

ผลลัพธ์จาก Predict set ถูกผสมผสานกับค่าคะแนนรีวิว (rating) ด้วยวิธีผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) โดยชุดข้อมูลที่ถูกบันทึกจะถูกแยกเป็น 21 ชุด ดังตาราง 2 โดยรูปแบบข้อมูลของข้อมูลแต่ละชุดเป็นดังภาพประกอบ 14 ซึ่งข้อมูลทุกชุดของทั้ง 21 ชุดนี้ แต่ละชุดจะประกอบไปด้วย 5,000 ผู้ใช้งานตามจำนวนของชุดข้อมูล Predict set

ตาราง 2 ชุดข้อมูลที่จะนำไปใช้กับระบบแนะนำ

ชุดข้อมูล	สมการผสมเชิงเส้นแบบถ่วงน้ำหนัก	ค่าน้ำหนัก (α)
1-7	$H=\alpha \times R+(1-\alpha) \times S$	0, 0.1, 0.3, 0.5, 0.7, 0.9, 1
8-15	$H=\alpha \times R+(1-\alpha) \times S \times \beta$	0, 0.1, 0.3, 0.5, 0.7, 0.9, 1
16-21	$H=\alpha \times R+(1-\alpha) \times \beta$	0, 0.1, 0.3, 0.5, 0.7, 0.9, 1

	USER	ITEM	RATING
0	AEPZQSZRITWTSW4WK5XZDXCZQFFA	B00CVEQXS6	1.0
1	AHQ5BQ4GYF7PYWCEEITKB5YJV72Q	B01ARGKUME	1.0
2	AH5V5TB6YGCPSGDIYKKAJEYK4W6A	B0B8XNR6BG	1.0
3	AG5ZLUUHVQ3WWYZCZZJOOWRVRZQ	B00DMCVZ1G	0.0
4	AFTIB2A2DA5757Z3PTTI5SVFBGYA	B007U8P35M	1.0

ภาพประกอบ 14 ตัวอย่างชุดข้อมูลที่ผ่านมาการทำ sentiment analysis และผสมข้อมูลกับค่าเรตติ้งสมการเชิงเส้นถ่วงน้ำหนักแล้ว พร้อมนำไปใช้กับระบบแนะนำ

3.4 การนำชุดข้อมูลที่ได้มาใช้กับระบบแนะนำ

ชุดข้อมูลที่ได้ถูกนำไปใช้ในระบบแนะนำ 3 รูปแบบ ได้แก่ ระบบแนะนำแบบการกรองข้อมูลโดยพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้ (User-based Collaborative Filtering), ระบบแนะนำแบบการกรองข้อมูลโดยพึ่งพาผู้ใช้ร่วมโดยอิงรายการ (Item-based Collaborative Filtering) และระบบแนะนำแบบผสม (Hybrid Collaborative Filtering) ที่รวม User-based และ Item-based เข้าด้วยกันผ่านวิธีการผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) โดยกำหนดค่าน้ำหนักที่ 0.1, 0.3, 0.5, 0.7 และ 0.9

การประเมินผลระบบแนะนำในงานวิจัยนี้ใช้ตัวชี้วัด Root Mean Square Error (RMSE), Mean Absolute Error (MAE), Precision@K และ Recall@K เพื่อวัดความแม่นยำและประสิทธิภาพของระบบแนะนำในแต่ละรูปแบบอย่างละเอียดและเปรียบเทียบผลลัพธ์เพื่อระบุแนวทางที่เหมาะสมที่สุดสำหรับการใช้งานในขั้นตอนการอภิปรายผลและสรุปการดำเนินงานต่อไป

บทที่ 4

ผลการดำเนินงานวิจัย

งานวิจัยนี้เป็นการวิจัยเชิงกรณีศึกษาของการนำผลลัพธ์จากข้อมูลการวิเคราะห์ความรู้สึกมาใช้เพิ่มประสิทธิภาพให้กับระบบแนะนำแบบกรองผู้ใช้รวม ผู้วิจัยได้ดำเนินการศึกษาตามกระบวนการของแต่ละวิธี เพื่อให้บรรลุวัตถุประสงค์ในการวิจัยที่ได้กำหนดไว้ ดังนี้

4.1 ผลลัพธ์ของแบบจำลองวิเคราะห์ความรู้สึก

จากการปรับปรุงพารามิเตอร์ของแบบจำลอง SieBERT ด้วยคำสั่งดังภาพประกอบ 15

```
# Define TrainingArguments for Trainer with the adjusted parameters
training_args = TrainingArguments(
    output_dir='./results',
    learning_rate=6e-5,
    num_train_epochs=20,
    per_device_train_batch_size=32,
    per_device_eval_batch_size=32,
    warmup_steps=400,
    weight_decay=0.07,
    save_strategy="epoch",
    evaluation_strategy="epoch",
    load_best_model_at_end=True,
    lr_scheduler_type='linear'
)

# Initialize Trainer
trainer = Trainer(
    model=siebert_model,
    args=training_args,
    train_dataset=tokenized_train_dataset,
    eval_dataset=tokenized_val_dataset,
    compute_metrics=compute_metrics,
    callbacks=[EarlyStoppingCallback(early_stopping_patience=5)]
)
```

ภาพประกอบ 15 คำสั่งที่ใช้ในการปรับปรุงพารามิเตอร์แบบจำลอง SieBERT

โดยพารามิเตอร์ที่ปรับมีดังตาราง 3

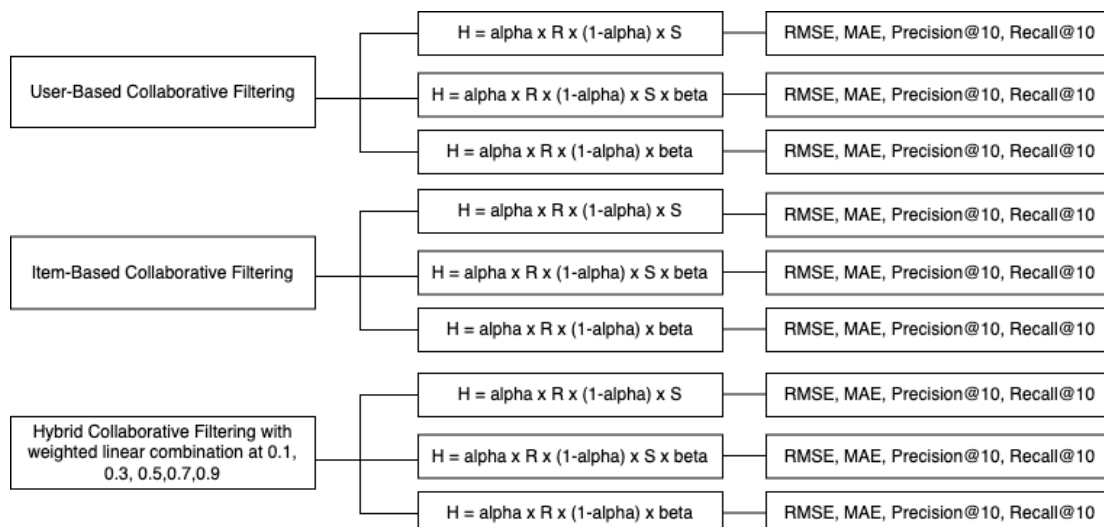
ตาราง 3 การปรับพารามิเตอร์

พารามิเตอร์	ค่าที่ตั้ง
Learning rate	6e-5
Num_train_epoch	20
Batch size	32
Warmup_steps	400
Weight_decay	0.07
Lr_scheduler_type	Linear
Early Stopping	5

ได้ผลลัพธ์จากชุดทดสอบเป็นค่า accuracy 92% ค่า AUC 69% เพิ่มขึ้นจากแบบจำลองวิเคราะห์ความรู้สึกแบบเทรนมาเรียบร้อยแล้ว (pretrained-model) ที่ให้ค่า accuracy 88% และค่า AUC 73% โดยค่า accuracy เพิ่มขึ้น 4.55% และค่า AUC ลดลง 5.48%

4.2 ผลลัพธ์ของการนำคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกไปผสมผสานกับคะแนนการจัดอันดับ(เรตติ้ง) และนำไปใช้ในระบบแนะนำ

ในการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับ (เรตติ้ง) คำนวณน้ำหนักของคะแนนแบ่งออกเป็นใช้คะแนนวิเคราะห์ความรู้สึกเท่านั้น ใช้คะแนนจากแบบจำลองวิเคราะห์ความรู้สึกผสมผสานกับคะแนนการจัดอันดับ(เรตติ้ง) ที่ค่าน้ำหนัก 0.1, 0.3, 0.5, 0.7, 0.9 และใช้คะแนนการจัดอันดับ(เรตติ้ง) เท่านั้น



ภาพประกอบ 16 การประเมินประสิทธิภาพของระบบแนะนำ

จากภาพประกอบ 16 ผู้วิจัยได้นำคะแนนที่ผสมกันในด้านน้ำหนักต่าง ๆ มาทดลองในระบบแนะนำ 3 แบบ ได้แก่ ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้ ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงรายการ และระบบแนะนำแบบผสม และใช้ตัวประเมินผล (evaluation metrics) เป็น RMSE, MAE, Precision@10 และ Recall@10

ทั้งนี้ ในแต่ละชนิดของระบบแนะนำ ได้ใช้สมการถ่วงน้ำหนัก 3 แบบ คือแบบที่ใช้ค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกผสมผสานกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times S$) แบบที่ใช้ค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและค่าความน่าจะเป็น (probability) ผสมผสานกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times S \times \beta$) และใช้เพียงค่าความน่าจะเป็นเพียงอย่างเดียวมาผสมผสานกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times \beta$) โดยที่ H คือ คะแนนที่ถูกผสมแล้ว α คือ ค่าน้ำหนัก β คือ ค่าความน่าจะเป็น (probability) ที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก S คือคะแนนที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก (คะแนน sentiment มีค่าเท่ากับ 0 หรือ 1) และ R คือ คะแนนการจัดอันดับ (เรตติ้ง)

ในการวิจัยนี้ได้ทำการวัดผลในสองส่วน คือ ส่วนการทำนายเรตติ้ง ใช้ค่าความคลาดเคลื่อนกำลังสองเฉลี่ย Root Mean Square Error (RMSE) และค่าความคลาดเคลื่อนสัมบูรณ์เฉลี่ย (Mean Absolute Error: MAE) เป็นมาตรวัดผล ส่วน top-N recommendation ใช้ precision@10 และ recall@10 เป็นมาตรวัด

4.2.1 ผลลัพธ์ของการใช้ค่าผลลัพธ์จากการวิเคราะห์ความรู้สึกมาผสานกับค่าคะแนนการจัดอันดับ (เรตติ้ง) เพื่อนำมาใช้กับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วมโดยอิงผู้ใช้ (User-based Collaborative Filtering)

การเปรียบเทียบประสิทธิภาพของการผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับ (เรตติ้ง) ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 1-5 ได้ผลการดำเนินการดังตาราง 4

ตาราง 4 การเปรียบเทียบประสิทธิภาพของการผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 1-5

#	Alpha	RMSE	MAE	Precision@10	Recall@10
H= $\alpha \times R + (1-\alpha) \times S$	0	0.0960	0.0286	0	0
	0.1	0.1378	0.0453	0	0
	0.3	0.2150	0.0734	0	0
	0.5	0.2933	0.1005	0.1120	0.9991
	0.7	0.3721	0.1275	0.1519	0.9988
	0.9	0.4512	0.1544	0.1572	0.9988
	1	0.4908	0.1678	0.1712	0.9986
H= $\alpha \times R + (1-\alpha) \times S \times \beta$	0	0.0888	0.0280	0	0.9890
	0.1	0.1304	0.0440	0	0.9890
	0.3	0.2093	0.0722	0	0.9890
	0.5	0.2892	0.0996	0	0.9890
	0.7	0.3697	0.1269	0.1499	0.9890
	0.9	0.4504	0.1542	0.1572	0.9890
	1	0.4908	0.1678	0.1712	0.9890

ตาราง 4 (ต่อ)

#	Alpha	RMSE	MAE	Precision@10	Recall@10
H= $\alpha \times R + (1-\alpha) \times \beta$	0	0.0963	0.0331	0	0
	0.1	0.1347	0.0462	0	0
	0.3	0.2131	0.0733	0	0
	0.5	0.2922	0.1004	0	0
	0.7	0.3716	0.1273	0.1543	0.9988
	0.9	0.4511	0.1543	0.1572	0.9988
	1	0.4908	0.1678	0.1712	0.9989

สรุปประสิทธิภาพของการผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงผู้ใช้จากตารางที่ 4 ได้ผลลัพธ์ดังตาราง 5

ตาราง 5 สรุปผลลัพธ์ของการผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติง 1-5

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	สมการถ่วงน้ำหนัก
RMSE ต่ำที่สุด	0.0888	0	$H = \alpha \times R + (1-\alpha) \times S \times \beta$
MAE ต่ำที่สุด	0.0280	0	$H = \alpha \times R + (1-\alpha) \times S \times \beta$
RMSE สูงที่สุด	0.4908	1	All
MAE สูงที่สุด	0.1678	1	All
PRECISION@10 สูงที่สุด	0.1712	1	All
RECALL@10 สูงที่สุด	0.9991	All	$H = \alpha \times R + (1-\alpha) \times S \times \beta$

สรุปประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 0-1 ได้ผลลัพธ์ดังตาราง 6

ตาราง 6 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวมโดยอิงผู้ใช้ โดยใช้คะแนนเรตติ้ง 0-1

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	สมการถ่วงน้ำหนัก
RMSE ต่ำที่สุด	0.1719	0.1	$H = \alpha \times R + (1 - \alpha) \times S \times \beta$
MAE ต่ำที่สุด	0.0955	0.9	$H = \alpha \times R + (1 - \alpha) \times \beta$
RMSE สูงที่สุด	0.1872	0	$H = \alpha \times R + (1 - \alpha) \times S$
MAE สูงที่สุด	0.1063	0	$H = \alpha \times R + (1 - \alpha) \times \beta$
PRECISION@10 สูงที่สุด	0.2731	0	$H = \alpha \times R + (1 - \alpha) \times S$
RECALL@10 สูงที่สุด	0.9895	0.3, 0.5, 0.7	$H = \alpha \times R + (1 - \alpha) \times S$

4.2.2 ผลลัพธ์ของการใช้ค่าผลลัพธ์จากการวิเคราะห์ความรู้สึกมาผสมผสานกับค่าคะแนนการจัดอันดับ (เรตติ้ง) เพื่อนำมาใช้กับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวมโดยอิงรายการ (Item-based Collaborative Filtering)

การเปรียบเทียบประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับ (เรตติ้ง) ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้รวมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 1-5 ได้ผลการดำเนินการดังตาราง 7

ตาราง 7 การเปรียบเทียบประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึก และคะแนนการจัดอันดับ (เรตติ้ง) ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้เข้าร่วมโดย อิงรายการ โดยใช้คะแนนเรตติ้ง 1-5

#	Alpha	RMSE	MAE	Precision@10	Recall@10
H= $\alpha \times R + (1-\alpha) \times S$	0	0.1033	0.0201	0	0
	0.1	0.2494	0.1197	0	0
	0.3	0.3230	0.1585	0	0
	0.5	0.4159	0.2001	0.1119	0.9990
	0.7	0.5179	0.2461	0.1519	0.9988
	0.9	0.6247	0.2938	0.1571	0.9988
	1	0.6792	0.3178	0.1712	0.9986
H= $\alpha \times R + (1-\alpha) \times S \times \beta$	0	0.1136	0.0531	0	0
	0.1	0.2357	0.1272	0	0
	0.3	0.3150	0.1633	0	0
	0.5	0.4116	0.2042	0	0
	0.7	0.5159	0.2488	0.1498	0.9988
	0.9	0.6242	0.2947	0.1571	0.9988
	1	0.6792	0.3178	0.1712	0.9986
H= $\alpha \times R + (1-\alpha) \times \beta$	0	0.1230	0.0604	0	0
	0.1	0.1644	0.0771	0	0
	0.3	0.2696	0.1264	0	0
	0.5	0.3839	0.1806	0	0
	0.7	0.5012	0.2354	0.1542	0.9988
	0.9	0.6196	0.2903	0.1571	0.9988
	1	0.6792	0.3178	0.1712	0.9986

จากตาราง 7 สรุปประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้รวมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 1-5 ได้ผลลัพธ์ดังตาราง 8

ตาราง 8 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้รวมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 1-5

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	สมการถ่วงน้ำหนัก
RMSE ต่ำที่สุด	0.1033	0	$H = \alpha \times R + (1-\alpha) \times S$
MAE ต่ำที่สุด	0.0201	0	$H = \alpha \times R + (1-\alpha) \times S$
RMSE สูงที่สุด	0.6792	1	All
MAE สูงที่สุด	0.3178	1	All
PRECISION@10 สูงที่สุด	0.1712	1	all
RECALL@10 สูงที่สุด	0.9990	0.5	$H = \alpha \times R + (1-\alpha) \times S$

สรุปประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้รวมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 0-1 ได้ผลลัพธ์ดังตาราง 9

ตาราง 9 สรุปผลลัพธ์ของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึงพาผู้ใช้รวมโดยอิงรายการ โดยใช้คะแนนเรตติ้ง 0-1

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	สมการถ่วงน้ำหนัก
RMSE ต่ำที่สุด	0.1807	0	$H = \alpha \times R + (1-\alpha) \times S \times \beta$
MAE ต่ำที่สุด	0.0598	1	all
RMSE สูงที่สุด	0.2758	0.9	$H = \alpha \times R + (1-\alpha) \times \beta$
MAE สูงที่สุด	0.1491	0.9	$H = \alpha \times R + (1-\alpha) \times \beta$

ตาราง 9 (ต่อ)

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	สมการถ่วงน้ำหนัก
PRECISION@10 สูงที่สุด	0.2730	0	$H = \alpha \times R + (1-\alpha) \times S$
RECALL@10 สูงที่สุด	0.9885	0.1	$H = \alpha \times R + (1-\alpha) \times S$

4.2.3 ผลลัพธ์ของการใช้ค่าผลลัพธ์จากการวิเคราะห์ความรู้สึกมาผสมกับค่าคะแนนการจัดอันดับ (เรตติ้ง) มาใช้ในระบบแนะนำแบบผสม

จากการนำค่าผลลัพธ์จากการวิเคราะห์ความรู้สึกมาผสมกับคะแนนการจัดอันดับเพื่อนำมาใช้ในระบบแนะนำแบบผสมระหว่างระบบแนะนำแบบกรองผู้ใช้รวมเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้ (User-based Collaborative Filtering) และระบบแนะนำแบบกรองผู้ใช้รวมเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงรายการ (Item-Based Collaborative Filtering) ด้วยวิธีผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) ที่ค่าถ่วงน้ำหนักผลลัพธ์จากการวิเคราะห์ความรู้สึกและค่าคะแนนการจัดอันดับเป็น 0, 0.1, 0.3, 0.5, 0.7, 0.9, 1 และค่าถ่วงน้ำหนักระหว่าง User-based และ Item-based (α_B) เป็น 0.1, 0.3, 0.5, 0.7, 0.9 โดยยิ่งค่าน้อย ยิ่งมีสัดส่วนของ Item-based มาก

$$3.1 H = \alpha \times R + (1-\alpha) \times S$$

$$3.2 H = \alpha \times R + (1-\alpha) \times S \times \beta$$

$$3.3 H = \alpha \times R + (1-\alpha) \times \beta$$

โดยที่ H คือ คะแนนที่ถูกผสมแล้ว α คือ ค่าน้ำหนัก β คือ ค่าความน่าจะเป็น (probability) ที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก S คือคะแนนที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก (คะแนน sentiment มีค่าเท่ากับ 0 หรือ 1) และ R คือ คะแนนการจัดอันดับ(เรตติ้ง)

ได้ผลการดำเนินการดังตาราง 10

ตาราง 10 การเปรียบเทียบประสิทธิภาพของการผสมผสานคะแนนจากแบบจำลองวิเคราะห์ความรู้สึก และคะแนนการจัดอันดับของระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนนเรตติ้ง 1-5

#	Alpha	Alpha_B	RMSE	MAE	Precision@10	Recall@10
H= alpha x R + (1-alpha) x S	0	0.1	0.0939	0.0209	0	0
		0.3	0.0790	0.0225	0	0
		0.5	0.0721	0.0242	0	0
		0.7	0.0752	0.0258	0	0
		0.9	0.0874	0.0274	0	0
	0.1	0.1	0.2259	0.1107	0	0
		0.3	0.1824	0.0945	0	0
		0.5	0.1469	0.0791	0	0
		0.7	0.1265	0.0642	0	0
		0.9	0.1284	0.0505	0	0
	0.3	0.1	0.2928	0.1467	0	0
		0.3	0.2391	0.1255	0	0
		0.5	0.1997	0.1071	0	0
		0.7	0.1842	0.0911	0	0
		0.9	0.1982	0.0781	0	0
	0.5	0.1	0.3772	0.1856	0.1119	0.9990
		0.3	0.3093	0.1595	0.1119	0.9990
		0.5	0.2617	0.1380	0.1119	0.9990
		0.7	0.2467	0.1199	0.1120	0.9991
		0.9	0.2697	0.1054	0.1120	0.9991
	0.7	0.1	0.4699	0.2285	0.1519	0.9988
		0.3	0.3858	0.1971	0.1519	0.9988
		0.5	0.3280	0.1713	0.1519	0.9988
		0.7	0.3113	0.1499	0.1519	0.9988
		0.9	0.3420	0.1330	0.1519	0.9988

ตาราง 10 (ต่อ)

#	Alpha	Alpha_B	RMSE	MAE	Precision@10	Recall@10
H= alpha x R + (1-alpha) x S (ต่อ)	0.9	0.1	0.5668	0.2734	0.1572	0.9988
		0.3	0.4656	0.2361	0.1572	0.9988
		0.5	0.3963	0.2056	0.1572	0.9988
		0.7	0.3770	0.1805	0.1572	0.9988
		0.9	0.4146	0.1609	0.1572	0.9988
	1	0.1	0.6162	0.2961	0.1712	0.9986
		0.3	0.5062	0.2558	0.1712	0.9986
		0.5	0.4310	0.2229	0.1712	0.9986
		0.7	0.4100	0.1958	0.1712	0.9986
		0.9	0.4510	0.1749	0.1712	0.9986
H= alpha x R + (1-alpha) x S x beta	0	0.1	0.1031	0.0491	0	0
		0.3	0.0851	0.0425	0	0
		0.5	0.0738	0.0373	0	0
		0.7	0.0724	0.0329	0	0
		0.9	0.0813	0.0292	0	0
	0.1	0.1	0.2135	0.1166	0	0
		0.3	0.1724	0.0976	0	0
		0.5	0.1389	0.0803	0	0
		0.7	0.1196	0.0641	0	0
		0.9	0.1215	0.0495	0	0
	0.3	0.1	0.2856	0.1503	0	0
		0.3	0.2332	0.1276	0	0
		0.5	0.1947	0.1080	0	0
		0.7	0.1794	0.0910	0	0
		0.9	0.1930	0.0771	0	0

ตาราง 10 (ต่อ)

#	Alpha	Alpha_B	RMSE	MAE	Precision@10	Recall@10
H= alpha x R + (1-alpha) x S x beta (ต่อ)	0.5	0.1	0.3733	0.1886	0	0
		0.3	0.3060	0.1616	0	0
		0.5	0.2588	0.1390	0	0
		0.7	0.2436	0.1200	0	0
		0.9	0.2661	0.1047	0	0
	0.7	0.1	0.4681	0.2305	0.1498	0.9988
		0.3	0.3842	0.1985	0.1498	0.9988
		0.5	0.3264	0.1721	0.1498	0.9988
		0.7	0.3096	0.1501	0.1498	0.9988
		0.9	0.3398	0.1327	0.1498	0.9988
	0.9	0.1	0.5663	0.2741	0.1571	0.9988
		0.3	0.4652	0.2366	0.1572	0.9988
		0.5	0.3959	0.2059	0.1572	0.9988
		0.7	0.3764	0.1806	0.1572	0.9988
		0.9	0.4139	0.1608	0.1572	0.9988
	1	0.1	0.6162	0.2961	0.1712	0.9986
		0.3	0.5062	0.2558	0.1712	0.9986
		0.5	0.4310	0.2229	0.1712	0.9986
		0.7	0.4100	0.1958	0.1712	0.9986
		0.9	0.4510	0.1749	0.1712	0.9986
H= alpha x R + (1-alpha) x beta	0	0.1	0.1117	0.0559	0	0
		0.3	0.0925	0.0485	0	0
		0.5	0.0805	0.0428	0	0
		0.7	0.0790	0.0381	0	0
		0.9	0.0883	0.0343	0	0

ตาราง 10 (ต่อ)

#	Alpha	Alpha_B	RMSE	MAE	Precision@10	Recall@10
H= alpha x R + (1-alpha) x beta (ต่อ)	0.1	0.1	0.1494	0.0713	0	0
		0.3	0.1242	0.0627	0	0
		0.5	0.1092	0.0565	0	0
		0.7	0.1089	0.0514	0	0
		0.9	0.1232	0.0471	0	0
	0.3	0.1	0.2448	0.1170	0	0
		0.3	0.2027	0.1025	0	0
		0.5	0.1766	0.0913	0	0
		0.7	0.1737	0.0822	0	0
		0.9	0.1952	0.0751	0	0
	0.5	0.1	0.3485	0.1676	0	0
		0.3	0.2876	0.1458	0	0
		0.5	0.2480	0.1286	0	0
		0.7	0.2406	0.1144	0	0
		0.9	0.2680	0.1036	0	0
	0.7	0.1	0.4548	0.2189	0.1542	0.9988
		0.3	0.3744	0.1897	0.1543	0.9988
		0.5	0.3208	0.1663	0.1543	0.9988
		0.7	0.3082	0.1470	0.1543	0.9988
		0.9	0.3411	0.1322	0.1543	0.9988
	0.9	0.1	0.5622	0.2703	0.1571	0.9988
		0.3	0.4622	0.2338	0.1572	0.9988
		0.5	0.3942	0.2040	0.1572	0.9988
		0.7	0.3760	0.1795	0.1572	0.9988
		0.9	0.4144	0.1606	0.1572	0.9988

ตาราง 10 (ต่อ)

#	Alpha	Alpha_B	RMSE	MAE	Precision@10	Recall@10
H= alpha x R + (1-alpha) x beta (ต่อ)	1	0.1	0.6162	0.2961	0.1571	0.9986
		0.3	0.5062	0.2558	0.1572	0.9986
		0.5	0.4310	0.2229	0.1572	0.9986
		0.7	0.4100	0.1958	0.1572	0.9986
		0.9	0.4510	0.1749	0.1572	0.9986

สรุปประสิทธิภาพของการผสมคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม จากตารางที่ 10 ได้ผลลัพธ์ดังตาราง 11 ดังนี้

ตาราง 11 สรุปผลลัพธ์ของการผสมคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนนเรตติ้ง 1-5

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	ค่า alpha_b	สมการถ่วงน้ำหนัก
RMSE ต่ำที่สุด	0.0721	0	0.5	H= alpha x R + (1-alpha) x S
MAE ต่ำที่สุด	0.0209	0	0.1	H= alpha x R + (1-alpha) x S
RMSE สูงที่สุด	0.4510	1	0.9	all
MAE สูงที่สุด	0.2961	1	0.1	all
PRECISION@10 สูงที่สุด	0.1572	1	all	all
RECALL@10 สูงที่สุด	0.9991	0.5	0.7 , 0.9	H = alpha x R + (1-alpha) x S

สรุปประสิทธิภาพของการผสมคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนนเรตติ้ง 0-1 ได้ผลลัพธ์ดังตาราง 12

ตาราง 12 สรุปผลลัพธ์ของการผสมคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและคะแนนการจัดอันดับในระบบแนะนำเทคนิคการกรองข้อมูลแบบผสม โดยใช้คะแนนเรตติ้ง 0-1

มาตรวัด	มีค่าเท่ากับ	ค่า alpha	ค่า alpha_b	สมการถ่วงน้ำหนัก
RMSE ต่ำที่สุด	0.1346	0	0.5	$H = \alpha \times R + (1-\alpha) \times S \times \beta$
MAE ต่ำที่สุด	0.0638	1	0.1	all
RMSE สูงที่สุด	0.2270	0.9	0.1	$H = \alpha \times R + (1-\alpha) \times S$
MAE สูงที่สุด	0.1431	0.9	0.1	$H = \alpha \times R + (1-\alpha) \times \beta$
PRECISION@10 สูงที่สุด	0.2731	0	0.3 , 0.5	$H = \alpha \times R + (1-\alpha) \times S$
RECALL@10 สูงที่สุด	0.9895	0.3	0.7 , 0.9	$H = \alpha \times R +$
		0.5	0.7 , 0.9	$(1-\alpha) \times S$
		0.7	0.7 , 0.9	

บทที่ 5

สรุปผลการวิจัย อภิปรายผล และข้อเสนอแนะ

ในการวิจัยเรื่องการเพิ่มศักยภาพระบบแนะนำด้วยการผสมผสานการวิเคราะห์ความรู้สึกด้วยแบบจำลองSieBERT เข้ากับระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วม สามารถแบ่งหัวข้อสรุปผลได้ดังต่อไปนี้

5.1 สรุปผลการวิจัย

5.2 อภิปรายผลการวิจัย

5.3 ข้อเสนอแนะ

5.1 สรุปผลการวิจัย

งานวิจัยนี้ศึกษาการนำค่าผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึก โดยใช้แบบจำลอง SieBERT มาผสมผสานกับคะแนนการจัดอันดับ (เรตติ้ง) เพื่อนำไปใช้กับระบบแนะนำเทคนิคการกรองร่วม (Collaborative Filtering) ทั้งนี้ ได้ศึกษาการปรับปรุงไฮเปอร์พารามิเตอร์ของแบบจำลอง SieBERT และทำการวิเคราะห์ความรู้สึกจากข้อมูลรีวิว เมื่อได้คะแนนจากการวิเคราะห์ความรู้สึกก็นำผลลัพธ์ที่ได้ไปผสมผสานกับคะแนนการจัดอันดับ (เรตติ้ง) ทั้งนี้ ได้ทำการทดลองผสมผสานคะแนนด้วยสมการถ่วงน้ำหนัก 3 แบบ คือแบบที่ใช้ค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกผสมผสานกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times S$) แบบที่ใช้ค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึก และค่าความน่าจะเป็น(probability) ผสมผสานกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times S \times \beta$) และใช้เพียงค่าความน่าจะเป็นเพียงอย่างเดียวผสมผสานกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times \beta$) โดยที่ H คือคะแนนที่ถูกผสมผสานแล้ว α คือ ค่าน้ำหนัก β คือ ค่าความน่าจะเป็น (probability) ที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก S คือ คะแนนที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก (คะแนน sentiment มีค่าเท่ากับ 0 หรือ 1) และ R คือ คะแนนการจัดอันดับ(เรตติ้ง) เพื่อทดสอบประสิทธิภาพการทำนายคะแนนการจัดอันดับระบบแนะนำ 3 ประเภท ได้แก่ ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วมโดยอิงผู้ใช้ ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วมโดยอิงรายการ และระบบแนะนำแบบผสมระหว่างระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้เข้าร่วมโดยอิงผู้ใช้และอิงรายการ ด้วยวิธีผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination) และทำการวัดผลด้วย RMSE และ MAE ในส่วนของการทำนายคะแนนการจัดอันดับ และวัดผลในส่วนของ top-N recommendation ด้วย precision@10 พบว่า ในระบบแนะนำทุกแบบ การใช้ค่าผลลัพธ์จากการวิเคราะห์ความรู้สึกในระบบ

แนะนำเพียงอย่างเดียวให้ค่าความผิดพลาด (RMSE และ MAE) ต่ำที่สุด และในส่วนของ top-N recommendation การไม่ใช้คะแนนจากแบบจำลองวิเคราะห์ความรู้สึกเข้ามาผสมเลย (ใช้แต่คะแนนการจัดอันดับ (เรตติ้ง) เพียงอย่างเดียว) ได้ผลลัพธ์ดีที่สุด หากเป็นการผสมคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกเข้ากับคะแนนการจัดอันดับ (เรตติ้ง) การใช้ค่าน้ำหนัก 0.1 ให้ผลลัพธ์ที่ดีที่สุดในทั้ง 3 ระบบแนะนำ และจากการวัดผลในส่วนของ top-N recommendation ด้วย precision@10 พบว่า การใช้ค่าคะแนนการจัดอันดับ (เรตติ้ง) เพียงอย่างเดียว ให้ผลลัพธ์ที่ดี ส่วนมาตรวัด recall@10 ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้ ผลลัพธ์ที่ดีเมื่อ $\alpha = 0.5$ โดยมีค่าเท่ากับ 0.9991 ที่สมการถ่วงน้ำหนักแบบ $H = \alpha \times R + (1-\alpha) \times S \times \beta$ ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงรายการ ค่า recall@10 สูงที่สุดเมื่อ $\alpha = 0.5$ โดยมีค่าเท่ากับ 0.9990 ที่สมการถ่วงน้ำหนักแบบ $H = \alpha \times R + (1-\alpha) \times S$ และในระบบแนะนำแบบผสม ค่า recall@10 สูงที่สุดเมื่อ $\alpha = 0.5$ และ $\alpha_b = 0.7$ และ 0.9 โดยมีค่าเท่ากับ 0.9991 ที่สมการถ่วงน้ำหนักแบบ $H = \alpha \times R + (1-\alpha) \times S$

5.2 อภิปรายผลการวิจัย

งานวิจัยนี้ศึกษาการนำค่าผลลัพธ์จากแบบจำลองวิเคราะห์ความรู้สึก โดยใช้แบบจำลอง SieBERT มาผสมกับคะแนนการจัดอันดับ (เรตติ้ง) 3 วิธี คือแบบที่ใช้ค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกผสมกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times S$) แบบที่ใช้ค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกและค่าความน่าจะเป็น (probability) ผสมกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times S \times \beta$) และใช้เพียงค่าความน่าจะเป็นเพียงอย่างเดียวผสมกับคะแนนการจัดอันดับ ($H = \alpha \times R + (1-\alpha) \times \beta$) โดยที่ H คือคะแนนที่ถูกผสมแล้ว α คือ ค่าน้ำหนัก β คือ ค่าความน่าจะเป็น (probability) ที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก S คือคะแนนที่ได้จากแบบจำลองวิเคราะห์ความรู้สึก (คะแนน sentiment มีค่าเท่ากับ 0 หรือ 1) และ R คือ คะแนนการจัดอันดับ (เรตติ้ง) เพื่อนำไปใช้กับระบบแนะนำเทคนิคการกรองร่วม (Collaborative Filtering) 3 แบบ คือ ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้ ระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงรายการ และระบบแนะนำแบบผสมระหว่างระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้ร่วมโดยอิงผู้ใช้และอิงรายการ ด้วยวิธีผสมเชิงเส้นแบบถ่วงน้ำหนัก (Weighted Linear Combination)

ในส่วนของการวิเคราะห์ความรู้สึก แบบจำลอง SieBERT สามารถใช้สำหรับการวิเคราะห์ความรู้สึกได้แม่นยำ โดยเฉพาะหลังปรับปรุงไฮเปอร์พารามิเตอร์ ให้ค่าความแม่นยำ (accuracy)

ถึง 92% เพิ่มขึ้น 4.55% เมื่อเทียบกับแบบจำลองที่ถูกเทรนมาแล้วโดยไม่ปรับแต่งเพิ่มเติมที่ให้ค่าความแม่นยำ 88%

ในส่วนของระบบแนะนำ จากการทดลองพบว่า การใช้เพียงค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึก (Sentiment Score) ซึ่งมีค่าเป็น 0 หรือ 1 ส่งผลให้ค่าความผิดพลาด (RMSE และ MAE) ต่ำที่สุดในทุกระบบแนะนำที่นำมาทดสอบ สะท้อนให้เห็นว่าการจัดอันดับรายการแนะนำโดยอิงจากคะแนนความรู้สึกที่ได้จากรีวิวของผู้ใช้งาน สามารถสะท้อนความพึงพอใจที่แท้จริงของผู้ใช้ได้ดีกว่าค่าคะแนนเรตติ้งที่ให้ไว้แบบตัวเลข (1-5) ผลลัพธ์นี้ชี้ให้เห็นว่า ข้อมูลเรตติ้งที่เป็นตัวเลขอาจมีความคลาดเคลื่อน (noise) หรือไม่สอดคล้องกับความรู้สึกที่แท้จริงของผู้ใช้ ขณะที่รีวิวข้อความซึ่งผ่านการวิเคราะห์เชิงความรู้สึกสามารถจับความต้องการและความพึงพอใจของผู้ใช้ได้แม่นยำกว่า ซึ่งความแม่นยำของแบบจำลองวิเคราะห์ความรู้สึกเองก็อาจมีส่วนช่วยเพิ่มประสิทธิภาพของระบบแนะนำให้ดียิ่งขึ้น

อย่างไรก็ตาม ข้อจำกัดของแนวทางนี้คือ การลดรูปข้อมูลให้อยู่ในรูปแบบไบนารี (Binary) อาจทำให้ระบบไม่สามารถสะท้อนความละเอียดของความพึงพอใจของผู้ใช้ที่ซับซ้อนขึ้นได้ นอกจากนี้ ค่าคะแนนความรู้สึกที่ได้จากรีวิวถือเป็น implicit rating ซึ่งอาจไม่สะท้อนความพึงพอใจที่แท้จริงของผู้ใช้ในทุกกรณี การผสมผสานค่าคะแนนนี้เข้ากับ explicit rating (ค่าคะแนนการจัดอันดับ (เรตติ้ง) ที่ผู้ใช้ให้โดยตรง) อาจช่วยให้ระบบแนะนำสะท้อนความเป็นจริงได้แม่นยำขึ้น แม้ว่าการใช้ explicit rating จะส่งผลให้ค่า RMSE และ MAE สูงขึ้น แต่ก็อาจให้ผลลัพธ์ที่สอดคล้องกับความพึงพอใจของผู้ใช้มากกว่าในเชิงปฏิบัติ

ในการผสมผสานค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึก (Sentiment Analysis) เข้ากับค่าคะแนนการจัดอันดับ (Rating) จากการทดลองพบว่า ช่วยลดค่าความผิดพลาดของการพยากรณ์ (Prediction Error) อย่างมีนัยสำคัญ โดยพิจารณาจากมาตรวัด Root Mean Square Error (RMSE) และ Mean Absolute Error (MAE) ซึ่งสะท้อนถึงประสิทธิภาพที่ดีขึ้นของระบบแนะนำ แสดงให้เห็นว่าการเพิ่มข้อมูลจากการวิเคราะห์ความคิดเห็นช่วยให้ระบบแนะนำสามารถสะท้อนแนวโน้มของผู้ใช้ได้ดีขึ้น โดยในทุกระบบแนะนำให้ผลดีที่สุดที่ค่าน้ำหนัก (alpha) เป็น 0.1 โดยเมื่อพิจารณาแต่ละเทคนิค พบว่า ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งานโดยอิงผู้ใช้ การผสมผสานด้วยสมการ $H = \alpha \times R + (1-\alpha) \times S \times \beta$ ให้ผลลัพธ์ที่ดีที่สุดในขณะที่ในระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งานโดยอิงรายการ และระบบแนะนำแบบผสมระหว่างระบบแนะนำเทคนิคการกรองข้อมูลแบบพึ่งพาผู้ใช้งานโดยอิงผู้ใช้และอิงรายการ ด้วยวิธีผสมเชิงเส้น

แบบถ่วงน้ำหนัก (Weighted Linear Combination) การผสมผสานคะแนนด้วยสมการ $H = \alpha \times R + (1-\alpha) \times S$ ให้ผลลัพธ์ที่ดีที่สุด

เมื่อประเมินประสิทธิภาพของระบบแนะนำในบริบทของ top-N recommendation โดยใช้ precision@10 เป็นมาตรวัด พบว่าการใช้เพียงค่าคะแนนการจัดอันดับ (เรตติ้ง) เป็นปัจจัยหลักในการแนะนำให้ผลลัพธ์ที่ดีที่สุด โดยมีค่า precision@10 ≈ 0.17 ซึ่งหมายความว่า หากแนะนำ 10 รายการ ผู้ใช้มีแนวโน้มที่จะพบรายการที่ตรงกับความต้องการเพียง 1.7 รายการ เท่านั้น โดยมีค่า recall@10 เท่ากับ 0.9991 ซึ่งหมายความว่า ระบบสามารถครอบคลุมรายการทั้งหมดที่ผู้ใช้น่าจะชอบได้เกือบครบถ้วน

อย่างไรก็ตาม เมื่อนำค่าคะแนนจากแบบจำลองวิเคราะห์ความรู้สึก (Sentiment Analysis) มาผสมผสานเข้ากับค่าคะแนนการจัดอันดับ พบว่าค่า precision@10 ลดลงเป็น 0 แสดงให้เห็นว่าการรวมข้อมูลความคิดเห็นของผู้ใช้เข้ากับคะแนนการจัดอันดับอาจไม่ได้ช่วยปรับปรุงความแม่นยำของการแนะนำในบริบทของ top-N recommendation และอาจส่งผลกระทบต่อประสิทธิภาพของระบบแนะนำ

ผลลัพธ์ดังกล่าวสะท้อนให้เห็นว่าการผสมผสานความรู้สึกอาจไม่ได้เหมาะสมกับทุกกรณีของระบบแนะนำ โดยเฉพาะเมื่อเป้าหมายคือการคัดกรองและแนะนำ N รายการที่ดีที่สุดให้กับผู้ใช้ เนื่องจากคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกอาจมีการกระจายที่แตกต่างจากคะแนนการจัดอันดับดั้งเดิม ส่งผลให้รายการที่มีคะแนนสูงสุดจากการรวมค่าคะแนนอาจไม่ได้สะท้อนความสนใจของผู้ใช้ได้อย่างแม่นยำ

ดังนั้น ข้อค้นพบนี้ชี้ให้เห็นถึงข้อจำกัดของการนำคะแนนจากแบบจำลองวิเคราะห์ความรู้สึกมาใช้ร่วมกับค่าคะแนนการจัดอันดับในบริบทของ top-N recommendation โดยอาจจำเป็นต้องพิจารณาปัจจัยเพิ่มเติม เช่น วิธีการถ่วงน้ำหนัก เพื่อให้การแนะนำยังคงมีประสิทธิภาพสูงสุดในบริบทที่แตกต่างกัน

5.3 ข้อเสนอแนะ

1. ทำการทดลองในโดเมนอื่น ๆ เพื่อประเมินประสิทธิภาพของระบบแนะนำในบริบทที่ต่างกัน โดยการทดลองในหลากหลายโดเมนจะช่วยให้เห็นภาพรวมของการทำงานของระบบแนะนำในเงื่อนไขที่หลากหลาย และอาจจะช่วยค้นพบแนวทางหรือปัจจัยใหม่ ๆ ที่สามารถปรับปรุงประสิทธิภาพได้

2. ศึกษาวิธีการเพิ่มข้อมูลบริบท (เช่น เมตาเดตาของรายการ หรือข้อมูลพฤติกรรมอื่น ๆ) โดยการเพิ่มข้อมูลเหล่านี้จะช่วยให้ระบบสามารถเข้าใจและทำนายความต้องการของผู้ใช้ได้ดียิ่งขึ้น

บรรณานุกรม

- Aggarwal, C. C. (2016). *Recommender Systems*. Switzerland: Springer Cham.
- Dang, C. N., Moreno-García, M. N., & Prieta, F. D. I. (2021). An Approach to Integrating Sentiment Analysis into Recommender Systems. *Sensors*, 21(16), 5666.
- Isinkaye, F. O., Folajimi, Y. O., & Ojokoh, B. A. (2015). Recommendation systems: Principles, methods and evaluation. *Egyptian Informatics Journal*, 16(3), 261-273.
- Karabila, I., Darraz, N., El-Ansari, A., Alami, N., & El Mallahi, M. (2023). Enhancing collaborative filtering-based recommender system using sentiment analysis. *Future Internet*, 15(7), 235.
- Karabila, I., Darraz, N., EL-Ansari, A., Alami, N., & EL Mallahi, M. (2024). BERT-enhanced sentiment analysis for personalized e-commerce recommendations. *Multimedia Tools and Applications*, 83(19), 56463-56488.
- Karabila, I., Darraz, N., El-Ansari, A., Alami, N., Lazaar, M., & El Mallahi, M. (2023). *Recommendation system using Deep Learning-based sentiment analysis*. Paper presented at the 2023 Sixth International Conference on Vocational Education and Electrical Engineering (ICVEE).
- Morzelona, R., & Batra, S. (2022). Sentiment-aware Content Recommendation using LSTM-based Collaborative Filtering. *Research Journal of Computer Systems and Engineering*, 3(2), 32-38.
- Patcharawongsakda, E. (2021). *A little book of Big data and machine learning*. Bangkok: Infopress Group.
- Sahu, S., Kumar, R., MohdShafi, P., Shafi, J., Kim, S., & Ijaz, M. F. (2022). A Hybrid Recommendation System of Upcoming Movies Using Sentiment Analysis of YouTube Trailer Reviews. *Mathematics*, 10(9), 1568.
- Sunwall, E. (2021). *Recognize Strategic Opportunities with Long-Tail Data*. Nielsen Norman Group. <https://www.nngroup.com/articles/long-tail/>

Upadhyaya, P. (2023). *Study of Mathematical Model for User-based Collaborative Filtering and Item-based Collaborative Filtering*. Research Gate. https://www.researchgate.net/publication/366902172_Study_of_Mathematical_Model_for_User-based_Collaborative_Filtering_and_Item-based_Collaborative_Filtering

