



การวิเคราะห์ความรู้สึกของรีวิวโรงแรมในเกาะกูดโดยใช้การเรียนรู้เชิงลึก

SENTIMENT ANALYSIS OF HOTEL REVIEWS ON KOH KOOD USING DEEP LEARNING
TECHNIQUES



สุชาติ ทองเกษร

บัณฑิตวิทยาลัย มหาวิทยาลัยศรีนครินทรวิโรฒ

2567

การวิเคราะห์ความรู้สึกรักของวีรวิธโรงแรมในเกาะภูเก็ตโดยใช้การเรียนรู้เชิงลึก



การค้นคว้าอิสระนี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร

วิทยาศาสตร์มหาบัณฑิต สาขาวิทยาการข้อมูล

คณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ

ปีการศึกษา 2567

ลิขสิทธิ์ของมหาวิทยาลัยศรีนครินทรวิโรฒ

SENTIMENT ANALYSIS OF HOTEL REVIEWS ON KOH KOOD USING DEEP LEARNING
TECHNIQUES



SUCHART TONGKESORN

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

การวิเคราะห์ความรู้สึกของรีวิวโรงแรมในเกาะภูเก็ตโดยใช้การเรียนรู้เชิงลึก

ของ

สุชาติ ทองเกษร

ได้รับอนุมัติจากบัณฑิตวิทยาลัยให้นับเป็นส่วนหนึ่งของการศึกษาตามหลักสูตร

วิทยาศาสตรมหาบัณฑิต สาขาวิทยาการข้อมูล

ของมหาวิทยาลัยศรีนครินทรวิโรฒ

(รองศาสตราจารย์ นายแพทย์จักรชัย เอกปัญญากุล)

คณบดีบัณฑิตวิทยาลัย

คณะกรรมการสอบปากเปล่าการค้นคว้าอิสระ

..... อาจารย์ที่ปรึกษา

(ผู้ช่วยศาสตราจารย์ ดร.วราภรณ์ วิทยานนท์)

..... ประธานกรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.จันตรี ผลประเสริฐ)

..... กรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.ศิริสรรพ เหล่าหะเกียรติ)

ชื่อเรื่อง	การวิเคราะห์ความรู้สึกของรีวิวโรงแรมในเกาะกูดโดยใช้การเรียนรู้เชิงลึก
ผู้วิจัย	สุชาติ ทองเกษร
ปริญญา	วิทยาศาสตร์มหาบัณฑิต
ปีการศึกษา	2567
อาจารย์ที่ปรึกษา	ผู้ช่วยศาสตราจารย์ ดร.วราภรณ์ วิทยานนท์

การเติบโตของอุตสาหกรรมการท่องเที่ยวในประเทศไทย โดยเฉพาะในพื้นที่เกาะกูด ส่งผลให้รีวิวโรงแรมออนไลน์กลายเป็นแหล่งข้อมูลที่มีคุณค่าสำหรับการประเมินความพึงพอใจของลูกค้า งานวิจัยนี้มีวัตถุประสงค์เพื่อวิเคราะห์ความรู้สึกที่แสดงออกในข้อความรีวิวโรงแรมจากเกาะกูด โดยใช้เทคนิคการเรียนรู้เชิงลึก (Deep Learning) เพื่อจำแนกความคิดเห็นออกเป็น 3 ประเภท ได้แก่ เชิงบวก เชิงกลาง และเชิงลบ การจำแนกในลักษณะสามระดับนี้ช่วยให้สามารถวิเคราะห์ความคิดเห็นของลูกค้าได้อย่างละเอียดมากยิ่งขึ้น เนื่องจากมักมีความรู้สึกที่หลากหลายอยู่ในข้อความเดียวกัน จากการศึกษาได้รวบรวมข้อความรีวิวจำนวน 3,000 รายการ และพัฒนาแบบจำลอง Deep Learning ทั้งหมด 5 แบบ ได้แก่ FastText + LSTM, FastText + BiLSTM, WangchanBERTa + LSTM, WangchanBERTa + BiLSTM และ FastText + GRU โดยได้ทำการปรับแต่งค่าพารามิเตอร์ (Hyperparameter) และประเมินประสิทธิภาพของแต่ละแบบจำลองโดยใช้เกณฑ์มาตรฐาน ได้แก่ ความแม่นยำ (Accuracy) ค่าความแม่นยำจำแนก (Precision) ค่าการดึงกลับ (Recall) และค่า F1-score WangchanBERTa ได้รับเลือกเป็นหนึ่งในเทคนิคการฝังความหมายของข้อความ (embedding technique) เนื่องจากได้รับการออกแบบมาโดยเฉพาะสำหรับภาษาไทย ซึ่งช่วยเพิ่มความสามารถในการจับบริบทและความหมายแฝงในข้อความภาษาไทยได้อย่างแม่นยำ ลักษณะเฉพาะนี้ทำให้ WangchanBERTa เหมาะอย่างยิ่งในการวิเคราะห์โครงสร้างที่ซับซ้อนของข้อความที่สร้างโดยผู้ใช้งาน ผลการทดลองแสดงให้เห็นว่าแบบจำลอง WangchanBERTa + LSTM ให้ประสิทธิภาพสูงสุด โดยมีความแม่นยำอยู่ที่ 69.40% และค่า F1-score อยู่ที่ 69.30% ข้อค้นพบนี้ชี้ให้เห็นว่า การใช้ embedding ที่ออกแบบมาเฉพาะสำหรับภาษาไทยร่วมกับ LSTM ให้ผลลัพธ์ดีกว่าแบบจำลองที่ใช้ embedding ทั่วไป งานวิจัยนี้จึงเน้นย้ำถึงประสิทธิภาพของการใช้ embedding เฉพาะภาษาในการเพิ่มความแม่นยำของแบบจำลอง และเสนอแนะให้มีการทดลองเพิ่มเติมกับชุดข้อมูลที่หลากหลายยิ่งขึ้น รวมทั้งศึกษาผลของเทคนิคการปรับสมดุลข้อมูลรูปแบบต่าง ๆ ที่อาจช่วยปรับปรุงประสิทธิภาพของการจำแนกความรู้สึก

คำสำคัญ : การวิเคราะห์ความรู้สึก, รีวิวโรงแรม, เกาะกูด, การเรียนรู้เชิงลึก
, WangchanBERTa, LSTM

Title	SENTIMENT ANALYSIS OF HOTEL REVIEWS ON KOH KOD USING DEEP LEARNING TECHNIQUES
Author	SUCHART TONGKESORN
Degree	MASTER OF SCIENCE
Academic Year	2024
Advisor	Assistant Professor Dr. Waraporn Viyanon

The growth of the tourism industry in Thailand, particularly in the Koh Kood area, has made online hotel reviews a valuable source of data for assessing customer satisfaction. This study aims to analyze the sentiments expressed in hotel review texts from Koh Kood using deep learning techniques to classify opinions into three categories: positive, neutral, and negative. This three-way classification enables a more nuanced analysis of customer feedback, which often contains mixed sentiments. A total of 3,000 review texts were collected, and five deep learning models were developed and compared: FastText + LSTM, FastText + BiLSTM, WangchanBERTa + LSTM, WangchanBERTa + BiLSTM, and FastText + GRU. Hyperparameter tuning was performed, and model performance was evaluated using standard metrics: accuracy, precision, recall, and F1-score. WangchanBERTa was chosen as one of the embedding techniques due to its design specifically for the Thai language, which enhances its ability to capture contextual nuances in Thai text. This characteristic makes it particularly suitable for analyzing the complex structure of user-generated reviews. Experimental results showed that the WangchanBERTa + LSTM model achieved the highest performance, with an accuracy of 69.40% and an F1-score of 69.30%. The findings suggest that Thai-specific contextual embeddings, combined with LSTM, outperform models that use general-purpose embeddings. The study highlights the effectiveness of employing language-specific embeddings to enhance model accuracy and recommends further experimentation with more diverse datasets, as well as an investigation into the impact of various data balancing techniques to improve sentiment classification performance.

Keyword: Sentiment Analysis, Hotel Reviews, Koh Kood, Deep Learning, WangchanBERTa, LSTM

กิตติกรรมประกาศ

ข้าพเจ้าขอกราบขอบพระคุณ ผู้ช่วยศาสตราจารย์ ดร.วราภรณ์ วิทยานนท์ อาจารย์ที่ปรึกษาโครงการวิทยานิพนธ์ จากภาควิชาวิทยาการข้อมูล คณะวิทยาศาสตร์ ที่ได้ให้คำแนะนำอย่างใกล้ชิด ให้ความช่วยเหลือในทุกขั้นตอนของการดำเนินงาน รวมถึงการให้กำลังใจอย่างสม่ำเสมอ จนทำให้ข้าพเจ้าสามารถดำเนินการวิจัยและจัดทำการค้นคว้าอิสระฉบับนี้จนสำเร็จลุล่วง

ข้าพเจ้าขอขอบคุณ คณะวิทยาศาสตร์ ภาควิชาวิทยาการคอมพิวเตอร์ ที่ให้การสนับสนุนทางวิชาการ สภาพแวดล้อมในการเรียนรู้ และโอกาสในการพัฒนาทักษะทางด้านวิทยาการข้อมูลอย่างครบถ้วน

นอกจากนี้ ข้าพเจ้าขอขอบคุณเพื่อนนักศึกษาร่วมรุ่นที่มีส่วนช่วยเหลือในการแลกเปลี่ยนความคิดเห็น และร่วมกันฝ่าฟันอุปสรรคตลอดระยะเวลาของการศึกษา ตลอดจนบุคคลที่มีส่วนช่วยให้การวิจัยครั้งนี้สำเร็จ

ท้ายที่สุด ข้าพเจ้าขอขอบคุณครอบครัวที่เป็นกำลังใจสำคัญในทุกช่วงเวลา ตลอดจนเป็นแรงผลักดันให้ข้าพเจ้ามีความมุ่งมั่นและประสบความสำเร็จ

ข้าพเจ้าขอแสดงความขอบคุณด้วยความซาบซึ้งใจอย่างยิ่ง

สุชาติ ทองเกษร

สารบัญ

หน้า

บทคัดย่อภาษาไทย.....	ง
บทคัดย่อภาษาอังกฤษ.....	จ
กิตติกรรมประกาศ	ฉ
สารบัญ	ช
สารบัญตาราง	ญ
สารบัญรูปภาพ.....	ฎ
บทที่ 1 บทนำ.....	1
1.1 ความเป็นมาของงานวิจัย.....	1
1.2 ความมุ่งหมายของงานวิจัย	2
1.3 ความสำคัญของงานวิจัย	2
1.4 ขอบเขตของงานวิจัย	3
1.5 สมมติฐานของงานวิจัย.....	3
1.6 ประโยชน์ที่คาดว่าจะได้รับ	4
บทที่ 2 ทบทวนวรรณกรรม	5
2.1 การวิเคราะห์ความรู้สึกร.....	6
2.2 การสร้างคุณลักษณะด้วยเทคนิค TF-IDF	6
2.3 การสร้างคุณลักษณะด้วยเทคนิค Word Embedding	8
2.4 ทฤษฎีเกี่ยวกับ LSTM (Long Short-Term Memory)	10
2.5 ทฤษฎีเกี่ยวกับ GRU (Gated Recurrent Unit).....	11

2.6 ทฤษฎีเกี่ยวกับ BiLSTM (Bidirectional Long Short-Term Memory).....	12
2.7 ทฤษฎีเกี่ยวกับ WangchanBERTa	14
2.8 ทฤษฎีเกี่ยวกับ FastText.....	15
2.9 ทฤษฎีเกี่ยวกับการประเมินวัดความแม่นยำของแต่ละอัลกอริทึม.....	17
2.10 งานวิจัยที่เกี่ยวข้อง.....	18
บทที่ 3 วิธีการดำเนินงาน.....	29
3.1 แผนการดำเนินงานวิจัย	29
3.2 การเก็บรวบรวมข้อมูลและการคัดเลือก.....	30
3.3 กระบวนการทำงานของแบบจำลอง	34
3.4 การเตรียมข้อมูล	36
3.5 การสร้างแบบจำลอง	39
3.6 การประเมินผลแบบจำลอง	44
บทที่ 4 ผลการศึกษา.....	46
4.1 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM.....	46
4.2 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM.....	48
4.3 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM	50
4.4 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM.....	52
4.5 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU	54
4.6 การเปรียบเทียบผลการวิเคราะห์ความรู้สึกระหว่าง LSTM และ BiLSTM ด้วย FastText และ WangchanBERTa.....	56
บทที่ 5 สรุปผล อภิปรายผล และข้อเสนอแนะ	59
5.1 สรุปผล	59
5.2 อภิปรายผล.....	60

5.3 ข้อเสนอแนะ.....	61
บรรณานุกรม.....	63



สารบัญตาราง

หน้า

ตาราง 1 ตัวอย่างข้อมูลผลการรีวิวการใช้บริการโรงแรม.....	7
ตาราง 2 การคำนวณ TF จากตัวอย่างข้อมูล.....	7
ตาราง 3 ผลการคำนวณ TF-IDF ของคำว่า “สะอาด”	8
ตาราง 4 ตัวอย่างข้อมูลเวกเตอร์ 3 มิติ.....	9
ตาราง 5 สรุปภาพรวมงานวิจัยที่เกี่ยวข้อง.....	25
ตาราง 6 แผนการดำเนินงาน	29
ตาราง 7 ค่าพารามิเตอร์เบื้องต้นของแบบจำลองทั้ง 5 เทคนิค	41
ตาราง 8 Fine-tuning Hyperparameters ของ FastText + LSTM / BiLSTM / GRU	42
ตาราง 9 Fine-tuning Hyperparameters ของ WangchanBERTa + LSTM / BiLSTM	43
ตาราง 10 ผลลัพธ์ของแต่ละแบบจำลองก่อนการปรับตัวแปร	56
ตาราง 11 ผลหลังจากการปรับพารามิเตอร์.....	57

สารบัญรูปภาพ

หน้า

ภาพประกอบ 1 โครงสร้างภายในของ Long Short-Term Memory (LSTM)	10
ภาพประกอบ 2 โครงสร้างภายในของ Gated Recurrent Unit (GRU)	12
ภาพประกอบ 3 สถาปัตยกรรมของ FastText ที่ใช้ในงาน text classification.....	15
ภาพประกอบ 4 หน้าจอการใช้งาน Instant Data Scraper.....	30
ภาพประกอบ 5 ตัวอย่างข้อมูลหลังจาก scraping data.....	31
ภาพประกอบ 6 ตัวอย่างข้อมูลผ่านการวิเคราะห์โครงสร้างข้อมูล (Structural Analysis).....	33
ภาพประกอบ 7 ตัวอย่างข้อมูลผ่านการวิเคราะห์ข้อความ (Textual Analysis) โดยแสดงผลในรูปแบบ Word cloud.....	34
ภาพประกอบ 8 ขั้นตอนการทำงานของแบบจำลอง.....	35
ภาพประกอบ 9 n-gram ในความคิดเห็นเชิงลบ.....	38
ภาพประกอบ 10 n-gram ในความคิดเห็นเป็นกลาง.....	38
ภาพประกอบ 11 n-gram ในความคิดเห็นเชิงบวก.....	39
ภาพประกอบ 12 ภาพประกอบสถาปัตยกรรม.....	44
ภาพประกอบ 13 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM ด้วยค่าคงที่.....	47
ภาพประกอบ 14 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM ด้วยการทำ Fine Tuning Hyperparameter	48
ภาพประกอบ 15 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM ด้วยค่าคงที่	49
ภาพประกอบ 16 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM ด้วยการทำ Fine Tuning Hyperparameter	50
ภาพประกอบ 17 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM ด้วยค่าคงที่	51

ภาพประกอบ 18 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM ด้วยการทำ Fine Tuning Hyperparameter.....	52
ภาพประกอบ 19 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM ด้วยค่าคงที่.....	53
ภาพประกอบ 20 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM ด้วยการทำ Fine Tuning Hyperparameter.....	54
ภาพประกอบ 21 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU ด้วยค่าปกติ.....	55
ภาพประกอบ 22 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU ด้วยการทำ Fine Tuning Hyperparameter.....	56



บทที่ 1

บทนำ

1.1 ความเป็นมาของงานวิจัย

ในสังคมปัจจุบัน โรงแรมและผู้ประกอบการในธุรกิจการบริการบนเกาะภูเก็ต ซึ่งเป็นแหล่งท่องเที่ยวที่มีชื่อเสียงและมีข้อมูลอยู่ใน Google Maps จำนวนมาก แพลตฟอร์มนี้ได้กลายมาเป็นหนึ่งในแหล่งข้อมูลสำคัญสำหรับผู้ที่ต้องการค้นหาและดูข้อมูลเกี่ยวกับสถานที่พักและบริการต่างๆ บนเกาะภูเก็ต โดยโรงแรม รีสอร์ท และที่พักที่ตั้งอยู่ในพื้นที่นี้มักได้รับการรีวิวและให้คะแนนจากผู้เข้าพักผ่านทาง Google Maps ซึ่งความคิดเห็นเหล่านี้มีบทบาทสำคัญ เนื่องจากส่งผลต่อชื่อเสียงของที่พักและการตัดสินใจเลือกใช้บริการของนักท่องเที่ยวที่กำลังมองหาที่พักบนเกาะภูเก็ต

ความคิดเห็นจากนักท่องเที่ยวที่เคยเข้าพักมักจะถูกแสดงออกในรูปแบบของภาษาธรรมชาติ (Natural Language) แต่อย่างไรก็ตาม คำที่ใช้ในการแสดงความคิดเห็นเหล่านี้อาจไม่ตรงตามหลักไวยากรณ์ และไม่มีกระบวนการอย่างชัดเจนว่าความคิดเห็นนั้นมีลักษณะเชิงบวก (Positive) เชิงลบ (Negative) หรือเป็นกลาง (Neutral) การวิเคราะห์ความคิดเห็นเหล่านี้จึงจำเป็นต้องใช้เทคโนโลยีการเรียนรู้ของเครื่อง (Machine Learning) เพื่อช่วยในการแยกแยะและวิเคราะห์ความรู้สึกที่แฝงอยู่ในความคิดเห็นอย่างมีประสิทธิภาพ

LSTM (Long Short-Term Memory) ซึ่งเป็นหนึ่งในสถาปัตยกรรมของ Recurrent Neural Network (RNN) และเป็นส่วนหนึ่งของการเรียนรู้เชิงลึก (Deep Learning) ได้รับการออกแบบมาเพื่อจดจำรูปแบบ (Patterns) ในระยะยาว ซึ่งเหมาะสมกับการวิเคราะห์ความคิดเห็นที่เป็นลำดับและต่อเนื่อง โดยสามารถเก็บข้อมูลความคิดเห็นจากอดีตและนำมาใช้ในการประมวลผลความคิดเห็นปัจจุบันได้อย่างแม่นยำและยืดหยุ่น

ด้วยเหตุนี้ ผู้วิจัยจึงได้ศึกษาและพัฒนาแบบจำลองสำหรับการวิเคราะห์ความรู้สึกโดยใช้เทคนิค LSTM เพื่อประเมินประสิทธิภาพของแบบจำลอง และนำมาใช้ในการวิเคราะห์ความคิดเห็นจากรีวิวภาษาไทยเกี่ยวกับโรงแรมและการบริการบนเกาะภูเก็ต เพื่อให้ผู้ประกอบการโรงแรมและรีสอร์ทในพื้นที่สามารถเข้าใจความต้องการและความรู้สึกของนักท่องเที่ยวได้ดียิ่งขึ้น ไม่ว่าจะเป็นความคิดเห็นเชิงบวกหรือเชิงลบ โดยนำข้อมูลเหล่านี้มาช่วยในการปรับปรุงบริการเพื่อเพิ่มความพึงพอใจของนักท่องเที่ยวที่มาเยือนเกาะภูเก็ต

1.2 ความมุ่งหมายของงานวิจัย

1. เพื่อศึกษาการพัฒนาแบบจำลอง RNN (Recurrent Neural Network) และเทคนิคที่พัฒนาขึ้นเพื่อแก้ปัญหาที่เกิดกับเทคนิคโครงข่ายประสาทเทียมแบบเวียนซ้ำ ได้แก่ เทคนิค LSTM (Long Short-Term Memory) , BiLSTM (Bidirectional Long Short-Term Memory) และเทคนิค GRU (Gated Recurrent Unit) เพื่อจดจำและวิเคราะห์รูปแบบของความคิดเห็นจากนักท่องเที่ยวที่ให้กับโรงแรมและที่พักบนเกาะภูเก็ตในรูปแบบของภาษาไทยที่มีความซับซ้อนและต่อเนื่อง

2. เพื่อทดสอบและประเมินความสามารถของแบบจำลอง LSTM , BiLSTM และ GRU ในการแยกแยะและประมวลผลความคิดเห็นที่มีลักษณะเชิงบวก เชิงลบ และเป็นกลางอย่างแม่นยำ

3. ปรับค่าพารามิเตอร์ (Fine-tuning Hyperparameters) เพื่อให้เทคนิคแต่ละเทคนิคที่มีประสิทธิภาพสูงสุด

1.3 ความสำคัญของงานวิจัย

ในยุคดิจิทัลปัจจุบัน ความคิดเห็นของลูกค้าที่แสดงออกผ่านช่องทางออนไลน์ เช่น Google Maps มีบทบาทสำคัญในการกำหนดทิศทางของธุรกิจโรงแรมและการบริการ โดยเฉพาะในพื้นที่แหล่งท่องเที่ยวอย่าง "เกาะภูเก็ต" ซึ่งเป็นจุดหมายปลายทางที่ได้รับความนิยมจากทั้งนักท่องเที่ยวชาวไทยและชาวต่างชาติ ความเห็นจากผู้เข้าพักเดิมไม่เพียงแต่เป็นเครื่องมือในการตัดสินใจของนักท่องเที่ยวรายใหม่ แต่ยังสะท้อนภาพลักษณ์และคุณภาพของบริการที่พักรได้อย่างชัดเจน

อย่างไรก็ตาม ข้อความรีวิวเหล่านี้มักแสดงออกด้วยภาษาธรรมชาติที่มีความหลากหลายทั้งในแง่ของโครงสร้างภาษา อารมณ์ และการใช้ถ้อยคำเฉพาะบุคคล ซึ่งการวิเคราะห์ด้วยวิธีดั้งเดิมอาจไม่สามารถแยกแยะความรู้สึกเชิงลึกที่ซ่อนอยู่ได้อย่างมีประสิทธิภาพ งานวิจัยนี้จึงมีความสำคัญในการนำเทคโนโลยีการเรียนรู้เชิงลึก (Deep Learning) โดยเฉพาะแบบจำลอง LSTM , BiLSTM และ WangchanBERTa ที่เหมาะกับการประมวลผลภาษาในบริบทของภาษาไทย มาใช้วิเคราะห์ความรู้สึกจากรีวิวของผู้ใช้บริการ

ผลลัพธ์ที่ได้จากงานวิจัยนี้จะช่วยให้ผู้ประกอบการที่พักและโรงแรมสามารถเข้าใจความรู้สึก ความคาดหวัง และปัญหาของลูกค้าได้อย่างแม่นยำและรวดเร็วยิ่งขึ้น สามารถนำไปสู่การปรับปรุงบริการในเชิงรุก เพิ่มประสบการณ์ที่ดีให้กับลูกค้า และสร้างความได้เปรียบทางการแข่งขันในตลาดที่มีการเปลี่ยนแปลงอยู่ตลอดเวลา นอกจากนี้ ข้อมูลเชิงลึกที่ได้ยังสามารถนำไปใช้ต่อยอดในเชิงกลยุทธ์ เช่น การวางแผนการตลาด การพัฒนาผลิตภัณฑ์ และการออกแบบประสบการณ์การเข้าพักในอนาคตอีกด้วย

1.4 ขอบเขตของงานวิจัย

1. ข้อมูลที่ใช้ในงานวิจัยนี้ได้จากเว็บไซต์ Google maps ซึ่งเป็นแพลตฟอร์มที่เปิดให้ผู้ใช้ทั่วไปสามารถให้คะแนนและแสดงความคิดเห็นเกี่ยวกับโรงแรมและที่พักได้ โดยจำกัดเฉพาะ โรงแรม และที่พักที่ตั้งอยู่ในพื้นที่เกาะภูเก็ต จังหวัดตราด เท่านั้น

2. ข้อมูลจะเริ่มตั้งแต่ 1 มกราคม 2023 จนถึง 31 มีนาคม 2024

3. ข้อมูลทั้งหมดจำนวน 3,000 ความคิดเห็น ได้จัดกลุ่มและแปะป้ายกำกับ (label) โดยแบ่งออกเป็น 3 หมวดหมู่ ได้แก่

- 1) ความคิดเห็นเชิงบวก (Positive) จำนวน 1,000 รายการ
- 2) ความคิดเห็นเป็นกลาง (Neutral) จำนวน 1,000 รายการ
- 3) ความคิดเห็นเชิงลบ (Negative) จำนวน 1,000 รายการ

4. แบบจำลองที่ใช้ศึกษามาทั้งหมด 3 ประเภทได้แก่ LSTM , BiLSTM และ GRU

5. การประเมินประสิทธิภาพของแบบจำลองจะใช้ Confusion Matrix เป็นเกณฑ์หลักในการวัดผลลัพธ์ของการจำแนกประเภทความคิดเห็นในแต่ละหมวดหมู่ โดยพิจารณาค่าตัวชี้วัดต่าง ๆ เช่น Accuracy, Precision, Recall และ F1-score

1.5 สมมติฐานของงานวิจัย

สมมติฐานของงานวิจัยครั้งนี้ ประกอบด้วย

1. แบบจำลอง LSTM (Long Short-Term Memory) เมื่อทำงานร่วมกับเทคนิคการฝังข้อความ (Embedding) เช่น FastText หรือ WangchanBERTa จะสามารถวิเคราะห์ข้อความภาษาไทยที่มีลักษณะซับซ้อนและมีบริบททางอารมณ์ที่เปลี่ยนแปลงภายในประโยคได้อย่างแม่นยำ เนื่องจาก LSTM มีความสามารถในการจดจำลำดับข้อมูลระยะยาว

2. แบบจำลอง BiLSTM (Bidirectional LSTM) เมื่อจับคู่กับโมเดลฝังข้อความเชิงบริบทอย่าง WangchanBERTa ซึ่งถูกฝึกจากข้อมูลภาษาไทยโดยเฉพาะ จะสามารถเข้าใจความหมายของข้อความที่มีการเปลี่ยนอารมณ์ระหว่างต้นประโยคและปลายประโยคได้ดีซึ่งกว่า LSTM ทิศทางเดียว โดยเฉพาะในกรณีที่ข้อความมีคำเชื่อมที่เปลี่ยนทิศทางของความรู้สึก เช่น “แต่” “แต่อย่างไรก็ตาม”

3. แบบจำลอง GRU (Gated Recurrent Unit) เมื่อใช้งานร่วมกับ FastText ซึ่งสามารถฝังคำศัพท์ได้อย่างมีประสิทธิภาพแม้ในกรณีที่มีคำที่ไม่ปรากฏในชุดข้อมูล (OOV – Out-of-Vocabulary) จะสามารถให้ผลการจำแนกความรู้สึกใกล้เคียงกับ LSTM ได้ แม้มีจำนวนพารามิเตอร์น้อยกว่าและใช้เวลาในการฝึกสั้นกว่า

4. แบบจำลองฝังข้อความ WangchanBERTa ซึ่งเป็นโมเดล Transformer ที่ได้รับการฝึกด้วยข้อมูลภาษาไทยจำนวนมาก จะสามารถแยกแยะความหมายในระดับบริบทได้ดีกว่า FastText ที่เป็นการฝังแบบไม่พิจารณาบริบท (non-contextualized) ส่งผลให้แบบจำลองที่ใช้ WangchanBERTa เป็น embedding layer มีแนวโน้มที่จะให้ผลลัพธ์ที่แม่นยำมากกว่า

1.6 ประโยชน์ที่คาดว่าจะได้รับ

1. ได้รับข้อมูลเชิงลึกจากความคิดเห็นของลูกค้า ซึ่งผ่านการจำแนกอารมณ์เชิงบวก เชิงลบ และเป็นกลางโดยแบบจำลองที่พัฒนา ทำให้ผู้ประกอบการโรงแรมและที่พักสามารถมองเห็น ภาพรวมคุณภาพของบริการ ได้อย่างชัดเจน และนำข้อมูลเหล่านี้ไปปรับปรุงส่วนที่บกพร่องหรือส่งเสริมจุดแข็งของตนเองได้ตรงจุดมากยิ่งขึ้น
2. สนับสนุนการวางแผนกลยุทธ์ทางการตลาดอย่างแม่นยำยิ่งขึ้น โดยข้อมูลจากการวิเคราะห์ความรู้สึกสามารถนำมาใช้ประกอบการตัดสินใจทางการตลาด เช่น การกำหนดกลุ่มเป้าหมาย การปรับปรุงข้อความโฆษณา หรือการสร้างประสบการณ์ที่ตรงใจลูกค้า
3. พัฒนาเครื่องมือวิเคราะห์ความคิดเห็นอัตโนมัติที่ช่วยลดเวลาและแรงงานในการประมวลผลข้อมูล ซึ่งช่วยให้ผู้ประกอบการสามารถเข้าถึงความคิดเห็นลูกค้าได้รวดเร็วและสรุปประเด็นสำคัญได้โดยไม่ต้องอ่านรีวิวทีละข้อความ
4. สร้างและเปรียบเทียบแบบจำลองการเรียนรู้เชิงลึก ได้แก่ LSTM, BiLSTM และ GRU ซึ่งจะช่วยให้ทราบว่าโมเดลใดเหมาะสมที่สุดสำหรับการวิเคราะห์ความคิดเห็นภาษาไทย โดยเฉพาะในบริบทของรีวิวโรงแรม และสามารถนำผลลัพธ์ไปใช้อ้างอิงหรือต่อยอดงานวิจัยอื่นได้

บทที่ 2

ทบทวนวรรณกรรม

ในการดำเนินงานวิจัยครั้งนี้ ผู้วิจัยได้ศึกษาเอกสารและงานวิจัยที่เกี่ยวข้องอย่างครอบคลุม เพื่อสร้างความเข้าใจเชิงลึกในองค์ความรู้ที่เกี่ยวข้องกับการวิเคราะห์ความรู้สึกจากข้อความภาษาไทย โดยใช้เทคนิคการเรียนรู้เชิงลึก และการประมวลผลภาษาธรรมชาติ ซึ่งเนื้อหาที่ศึกษาแบ่งออกเป็นประเด็นสำคัญดังต่อไปนี้ เริ่มจากการทำความเข้าใจแนวคิดพื้นฐานเกี่ยวกับ การวิเคราะห์ความรู้สึก (Sentiment Analysis) ซึ่งเป็นกระบวนการวิเคราะห์อารมณ์หรือความรู้สึกจากข้อความที่มนุษย์เขียนขึ้น เช่น ความคิดเห็นจากผู้ให้บริการหรือรีวิวออนไลน์ เพื่อจำแนกออกเป็นกลุ่ม เช่น เชิงบวก เชิงลบ หรือ เป็นกลาง ต่อมาได้ศึกษาวิธีการสร้างคุณลักษณะ (Feature Extraction) ที่มีบทบาทสำคัญต่อประสิทธิภาพของแบบจำลอง โดยพิจารณาทั้งเทคนิคเชิงสถิติอย่าง TF-IDF ซึ่งใช้ในการแปลงข้อความให้อยู่ในรูปเวกเตอร์ของความถี่เชิงถ่วงน้ำหนัก และเทคนิค Word Embedding ซึ่งใช้การเรียนรู้การแทนค่าคำศัพท์เป็นเวกเตอร์ในมิติต่าง ๆ ที่สามารถสะท้อนบริบทของคำได้อย่างมีประสิทธิภาพ ในด้านของแบบจำลองลำดับ (Sequence Models) ผู้วิจัยได้ศึกษาทฤษฎีที่เกี่ยวข้องกับ LSTM (Long Short-Term Memory) ซึ่งมีความสามารถในการจดจำข้อมูลในลำดับที่ยาว GRU (Gated Recurrent Unit) ซึ่งเป็นแบบจำลองที่ลดความซับซ้อนลงจาก LSTM แต่ยังสามารถรักษาข้อมูลสำคัญไว้ได้ และ BiLSTM (Bidirectional LSTM) ซึ่งสามารถเรียนรู้ข้อมูลจากทั้งทิศทางซ้ายไปขวาและขวาไปซ้าย ทำให้เข้าใจบริบทได้ครบถ้วนยิ่งขึ้น นอกจากนี้ ผู้วิจัยยังให้ความสำคัญกับการศึกษา โมเดลฝังข้อความ (Embedding Models) ได้แก่ FastText ซึ่งเป็นเทคนิคของ Facebook ที่สามารถจัดการกับคำที่ไม่เคยพบมาก่อน (OOV) ได้ดี และ WangchanBERTa ซึ่งเป็นโมเดล Transformer ที่ถูกฝึกมาเฉพาะกับข้อมูลภาษาไทย ช่วยให้การวิเคราะห์อารมณ์จากข้อความไทยมีความแม่นยำยิ่งขึ้น เพื่อประเมินผลของแต่ละอัลกอริทึมที่ใช้ในการทดลอง ผู้วิจัยได้ศึกษาวิธีการ วัดผลด้านประสิทธิภาพของแบบจำลอง เช่น Accuracy, Precision, Recall และ F1-Score ซึ่งเป็นเกณฑ์มาตรฐานในการวิเคราะห์เปรียบเทียบ สุดท้าย ผู้วิจัยได้พิจารณา งานวิจัยที่เกี่ยวข้อง ทั้งในประเทศและต่างประเทศ เพื่อศึกษาวิธีการที่นักวิจัยท่านอื่นใช้ในการแก้ปัญหาที่คล้ายคลึงกัน และใช้เป็นแนวทางในการออกแบบงานวิจัยฉบับนี้ให้เหมาะสม

2.1 การวิเคราะห์ความรู้สึก

ปัจจุบัน การวิเคราะห์ข้อมูลไม่ได้จำกัดเฉพาะข้อมูลเชิงตัวเลข (Structured Data) เท่านั้น แต่ยังครอบคลุมถึงข้อมูลที่ไม่มีโครงสร้าง (Unstructured Data) เช่น ข้อความ รูปภาพ และเสียง ซึ่งสามารถสะท้อนข้อมูลเชิงลึกเกี่ยวกับประสบการณ์ ความรู้สึก และความคิดเห็นของผู้ใช้งานได้อย่างมีคุณค่า

งานวิจัยนี้มุ่งเน้นการวิเคราะห์ความคิดเห็นในรูปแบบข้อความของผู้ใช้บริการโรงแรมบนเกาะกูด ซึ่งแสดงออกผ่านแพลตฟอร์มออนไลน์ โดยเฉพาะ Google Maps เพื่อทำความเข้าใจความรู้สึกที่แฝงอยู่ในข้อความเหล่านั้น โดยการวิเคราะห์ความรู้สึก (Sentiment Analysis) ด้วยการจำแนกความคิดเห็นออกเป็น 3 ประเภท ได้แก่ 1. ความคิดเห็นเชิงบวก (Positive) 2. ความคิดเห็นเป็นกลาง (Neutral) 3. ความคิดเห็นเชิงลบ (Negative)

การจำแนกความคิดเห็นดำเนินการโดยอาศัยเทคนิค การเรียนรู้ของเครื่อง (Machine Learning) และการเรียนรู้เชิงลึก (Deep Learning) ที่ฝึกฝนแบบจำลองจากชุดข้อมูลที่มีการระบุอารมณ์ไว้ล่วงหน้า (Supervised Learning) และนำแบบจำลองนั้นไปใช้กับชุดข้อมูลทดสอบ เพื่อทำนายแนวโน้มความรู้สึกของผู้ใช้บริการรายใหม่

การวิเคราะห์ความรู้สึกในลักษณะนี้ ช่วยให้ผู้ธุรกิจที่พักและโรงแรมสามารถเข้าใจความคิดเห็นของลูกค้าได้อย่างเป็นระบบและแม่นยำ นำไปสู่การปรับปรุงคุณภาพของบริการ การแก้ไขจุดอ่อน และการเสริมสร้างจุดแข็งให้สอดคล้องกับความคาดหวังของลูกค้า อันจะส่งผลต่อความพึงพอใจและการกลับมาใช้บริการในอนาคต

2.2 การสร้างคุณลักษณะด้วยเทคนิค TF-IDF

TF-IDF (Term Frequency-Inverse Document Frequency) เป็นเทคนิคพื้นฐานที่นิยมใช้ในการแปลงข้อความ (Text) ให้อยู่ในรูปแบบข้อมูลเชิงปริมาณ (Numerical Representation) เพื่อให้สามารถนำไปประมวลผลด้วยอัลกอริธึมการเรียนรู้ของเครื่อง (Machine Learning) ได้ โดย TF-IDF ช่วยสะท้อนถึงความสำคัญของคำแต่ละคำภายในเอกสารหนึ่ง ๆ เทียบกับชุดเอกสารทั้งหมด TF-IDF คำนวณความสำคัญของแต่ละคำในเอกสารโดยแบ่งออกเป็นสองส่วน ได้แก่

1. Term Frequency (TF) วัดจำนวนครั้งที่คำหนึ่ง ๆ ปรากฏในเอกสาร ยิ่งค่าปรากฏบ่อยแสดงว่าคำนั้นอาจมีความสำคัญในเนื้อหาของเอกสารนั้น โดยคำนวณดังสมการ (1)

$$TF(t) = \frac{\text{จำนวนครั้งที่คำ } t \text{ ปรากฏในเอกสาร}}{\text{จำนวนครั้งทั้งหมดในเอกสาร}} \quad (1)$$

ตาราง 1 ตัวอย่างข้อมูลผลการรีวิวการใช้บริการโรงแรม

Document	Review Text
D1	ห้องพัก สะอาด และ เงียบสงบ
D2	ห้องพัก กว้างขวาง แต่ ไม่สะอาด
D3	ทำเล สะดวก เดินทาง ง่าย และ สะอาด

ตาราง 2 การคำนวณ TF จากตัวอย่างข้อมูล

Document	จำนวนคำทั้งหมด	จำนวนคำว่า “สะอาด”	TF
D1	4	1	$1/4 = 0.25$
D2	5	1	$1/5 = 0.20$
D3	5	1	$1/5 = 0.20$

สมมติให้ตัวอย่างข้อมูลผลการรีวิวการใช้บริการโรงแรมจำนวน 3 ข้อความ ดังตาราง 1 โดยแสดงตัวอย่างการคำนวณค่าความถี่เชิงเทอม (Term Frequency: TF) ของคำว่า “สะอาด” ในเอกสาร ตัวอย่างจำนวน 3 ฉบับ ได้แก่ D1, D2 และ D3 โดย TF เป็นค่าที่ใช้วัดความถี่ของคำใดคำหนึ่งที่ปรากฏในเอกสารเทียบกับจำนวนคำทั้งหมดในเอกสารนั้น เพื่อสะท้อนความสำคัญเชิงสัมพัทธ์ของคำนั้นภายในเอกสาร (ตาราง 2)

ในเอกสาร D1 มีคำทั้งหมด 4 คำ โดยคำว่า “สะอาด” ปรากฏ 1 ครั้ง ทำให้ค่าความถี่เชิงเทอม (TF) เท่ากับ $1/4$ หรือ 0.25 เอกสาร D2 มีความยาว 5 คำ และมีคำว่า “สะอาด” 1 ครั้ง ทำให้ TF เท่ากับ $1/5$ หรือ 0.20 เช่นเดียวกันกับ D3 ซึ่งมีจำนวนคำทั้งหมด 5 คำ และมีคำว่า “สะอาด” 1 ครั้ง จึงได้ TF เท่ากับ $1/5$ หรือ 0.20

จากผลการคำนวณนี้สามารถสังเกตได้ว่า แม้คำว่า “สะอาด” จะปรากฏเพียงครั้งเดียวในทุกเอกสาร แต่ TF ของเอกสารแต่ละฉบับแตกต่างกันตามจำนวนคำรวมของเอกสารนั้น ซึ่งแสดงให้เห็นว่า TF สามารถสะท้อน “ความหนาแน่น” ของคำในแต่ละเอกสารได้อย่างชัดเจน

2. Inverse Document Frequency (IDF) ใช้เพื่อลดความสำคัญของคำที่พบในหลายเอกสาร โดยจะให้ค่าน้อยกับคำที่พบในหลายเอกสารและให้ค่าสูงขึ้นกับคำที่พบน้อยลง เพื่อช่วยเน้นคำที่มีลักษณะเฉพาะเจาะจงมากกว่า โดยคำนวณดังสมการ (2) เพื่อป้องกันค่าเป็นศูนย์

$$IDF(t) = \log \left(\frac{\text{จำนวนเอกสารทั้งหมด}}{1 + \text{จำนวนเอกสารที่มีคำ } t \text{ ปรากฏ}} \right) \quad (2)$$

การคำนวณ IDF ของคำว่า “สะอาด”

- จำนวนเอกสารทั้งหมด = 3

- จำนวนเอกสารที่มีคำว่า “สะอาด” = 3

ดังนั้น $IDF = \log(3 / (1 + 3)) = \log(0.75) \approx -0.125$

การคำนวณ TF-IDF ของคำว่า “สะอาด”

TF-IDF ของคำแต่ละคำในเอกสารสามารถคำนวณได้จากการนำค่า TF และ IDF มาคูณกันดังสมการ (3)

$$TF - IDF(t) = TF(t) \times IDF(t) \quad (3)$$

ตาราง 3 ผลการคำนวณ TF-IDF ของคำว่า “สะอาด”

เอกสาร	TF	IDF ≈ -0.125	TF-IDF
D1	0.25	-0.125	-0.031
D2	0.2	-0.125	-0.025
D3	0.2	-0.125	-0.025

จากตารางที่ 3 เมื่อได้ค่า TF-IDF ของแต่ละคำในเอกสารแล้ว สามารถนำค่าที่ได้ไปสร้างเป็น เวกเตอร์คุณลักษณะ (Feature Vector) เพื่อใช้ในการวิเคราะห์ข้อความต่อไป ซึ่งจากตัวอย่างจะเห็นว่าแม้คำว่า “สะอาด” จะปรากฏในทุกเอกสาร แต่ค่าความสำคัญโดยรวมกลับน้อย เพราะไม่สามารถ แยกแยะเอกสารได้

2.3 การสร้างคุณลักษณะด้วยเทคนิค Word Embedding

การวิเคราะห์ข้อความด้วยเทคนิคการเรียนรู้ของเครื่อง จำเป็นต้องแปลงข้อความที่เป็นภาษามนุษย์ให้อยู่ในรูปแบบที่เครื่องสามารถประมวลผลได้ หนึ่งในเทคนิคที่ได้รับความนิยมและมีประสิทธิภาพสูง คือ Word Embedding (Pennington et al., 2014) ซึ่งเป็นวิธีการแปลง

คำให้เป็นเวกเตอร์ตัวเลขที่สามารถสะท้อนความหมายและบริบทของคำนั้น ๆ ได้อย่างลึกซึ้งกว่าการใช้เทคนิคแบบดั้งเดิม เช่น Bag of Words หรือ TF-IDF

Word Embedding คือการแทนคำแต่ละคำในข้อความด้วยเวกเตอร์ที่มีมิติต่ำ (เช่น 100, 200, หรือ 300 มิติ) โดยเวกเตอร์แต่ละตัวจะถูกฝึกมาจากข้อมูลขนาดใหญ่ เพื่อให้คำที่มีความหมายใกล้เคียงกันมีเวกเตอร์ที่อยู่ใกล้กันในพื้นที่เวกเตอร์ (vector space)

ต่างจาก TF-IDF ซึ่งแสดงเพียงความถี่ของคำในเอกสาร Word Embedding สามารถเรียนรู้ความสัมพันธ์เชิงความหมายได้ เช่น คำว่า “โรงแรม” และ “ที่พัก” อาจอยู่ใกล้กันในเวกเตอร์ คำว่า “สะอาด” และ “สกปรก” จะอยู่คนละทิศทาง เนื่องจากมีความหมายตรงข้าม

ตาราง 4 ตัวอย่างข้อมูลเวกเตอร์ 3 มิติ

คำ	เวกเตอร์ตัวอย่าง (3 มิติ)
ห้องพัก	[0.25, 0.12, -0.31]
กว้างขวาง	[0.40, 0.08, -0.10]
และ	[0.01, 0.01, 0.00]
สะอาด	[0.60, 0.15, -0.20]
มาก	[0.10, 0.03, 0.00]

ตัวอย่างการใช้ WORD EMBEDDING สมมุติข้อความรีวิว่า “ห้องพักกว้างขวางและสะอาดมาก” ซึ่งหลังจากผ่านกระบวนการ Embedding แต่ละคำอาจถูกแปลงเป็นเวกเตอร์ดังตารางที่ 4 (สมมุติเป็นเวกเตอร์ขนาด 3 มิติ)

ประเภทของ Word Embedding ที่นิยมได้แก่

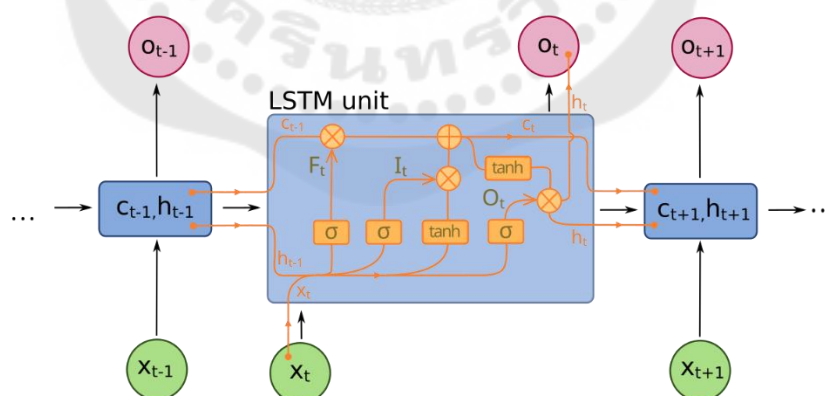
1. Word2Vec – พัฒนาโดย Google มี 2 โหมด: CBOW และ Skip-Gram ใช้การเรียนรู้จากบริบทของคำรอบข้าง
2. GloVe – พัฒนาโดย Stanford ใช้ข้อมูลการนับคำร่วม (co-occurrence matrix) และมีความแม่นยำในงานเชิงความหมาย
3. FastText – พัฒนาโดย Facebook ใช้ subword (เช่น พยางค์) ทำให้เหมาะกับภาษาอย่างภาษาไทยที่ไม่มีการเว้นวรรค

4. BPEmb / WangchanBERTa Embedding – รองรับภาษาไทยโดยตรง และเรียนรู้บริบทได้ดีกว่าในระดับ subword และ contextual

ดังนั้น Word Embedding เป็นเทคนิคที่สำคัญในการสร้างคุณลักษณะ (Feature Engineering) สำหรับงานประมวลผลภาษาธรรมชาติ (NLP) โดยเฉพาะเมื่อใช้ร่วมกับแบบจำลองเชิงลำดับ (Sequential Models) เช่น LSTM หรือ Transformer ซึ่งสามารถนำไปวิเคราะห์ความรู้สึกของข้อความได้อย่างมีประสิทธิภาพ ทั้งในภาษาอังกฤษและภาษาไทย โดยมีเครื่องมือและโมเดลที่หลากหลายให้เลือกใช้งาน

2.4 ทฤษฎีเกี่ยวกับ LSTM (Long Short-Term Memory)

LSTM เป็นเทคนิคที่ถูกพัฒนาขึ้นและนำเสนอในปี ค.ศ. 1997 โดย Hochreite และ Jurgen Schmidhuber (Hochreiter & Schmidhuber, 1997) เพื่อปรับปรุงข้อจำกัดที่เกิดขึ้นในขั้นตอนการฝึกการเรียนรู้ (Training) ของเทคนิค RNNs คือ การหายไปของเกรเดียนต์ (Gradient Vanishing) และการระเบิดของเกรเดียนต์ (Gradient Explosion) ของแบบจำลอง ด้วยวิธีถอยกลับ (back-propagation) เพื่อหาค่าพารามิเตอร์ (wikipedia, 2025) ได้แก่ น้ำหนัก (Weights) และไบแอส (Biases) เพื่อนำข้อผิดพลาดที่เกิดขึ้นให้แบบจำลองเรียนรู้ เพื่อลดค่าฟังก์ชันความสูญเสีย (Loss Function) ที่เทียบผลลัพธ์ที่ได้จากการพยากรณ์กับค่าจริง ดังภาพประกอบ 1 แสดงโครงสร้างภายในของ Long Short-Term Memory (LSTM)



ภาพประกอบ 1 โครงสร้างภายในของ Long Short-Term Memory (LSTM)

ที่มา : (Wikipedia, 2024b)

จากภาพประกอบ 1 ภายในโหนดของ LSTM ประกอบด้วย

1. X_t ข้อมูลอินพุตในช่วงเวลา t อินพุตในแต่ละช่วงเวลาอาจเป็นลำดับของคำ หรือค่าคุณลักษณะเชิงเวลา ซึ่งเป็นข้อมูลหลักที่โมเดลนำมาเรียนรู้ร่วมกับสถานะก่อนหน้า
2. C_t (Cell State) เป็นเส้นทางหลักที่ข้อมูลสามารถไหลผ่านได้โดยไม่สูญหาย โดยสามารถเลือกที่จะลบหรือเพิ่มข้อมูลใหม่ลงไปได้ในแต่ละช่วงเวลา ผ่านกลไกของประตูต่างๆ มีการอัปเดตค่าดังสมการ (4)

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \bar{C}_t \quad (4)$$

3. h_t (Hidden State) เป็นค่าที่ถูกส่งต่อไปยัง timestep ถัดไปและ/หรือใช้ในการพยากรณ์ในช่วงเวลานั้น คำนวณจากสมการ (5)

$$h_t = o_t \cdot \tanh(C_t) \quad (5)$$

4. f_t (Forget Gate) มีหน้าที่ในการตัดสินใจว่าข้อมูลใดจากสถานะเซลล์ก่อนหน้า C_{t-1} ควรถูกลบทิ้งโดย sigmoid activation function ดังสมการ (6)

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (6)$$

5. $i_t, \bar{C}_t$ (Input Gate) ใช้ในการควบคุมว่าข้อมูลใหม่จะถูกเพิ่มเข้าไปใน cell state หรือไม่ โดยประกอบด้วย 2 ส่วน

- 1) ตัวกรองการรับข้อมูล ดังสมการ (7)

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (7)$$

- 2) ค่าข้อมูลใหม่ที่เตรียมเข้าสู่ cell ดังสมการ (8)

$$\bar{C}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \quad (8)$$

6. o_t (Output Gate) ทำหน้าที่ควบคุมว่า cell state ที่อัปเดตแล้วควรถูกส่งออกมาเป็น hidden state หรือไม่ โดยคำนวณจากสมการ (9)

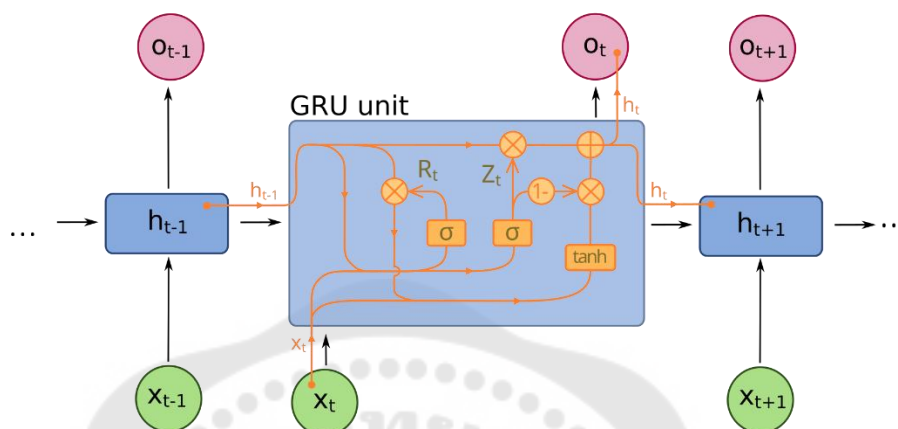
$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (9)$$

2.5 ทฤษฎีเกี่ยวกับ GRU (Gated Recurrent Unit)

GRU เป็นหนึ่งในรูปแบบของ RNN ที่ออกแบบมาเพื่อแก้ปัญหการสูญเสียมความจำในลำดับข้อมูลยาว (Long-term Dependencies) เช่นเดียวกับ LSTM แต่มีโครงสร้างที่เรียบง่ายกว่า จากภาพประกอบ 2 GRU มีโครงสร้าง gate หลักประกอบด้วย 2 ส่วน (Dey & Salem, 2017) ดังนี้

1. Update Gate เป็นประตูที่ควบคุมข้อมูลในปัจจุบันว่าควรถูกบันทึกไว้หรือแทนที่ข้อมูลเก่า

2. Reset Gate เป็นประตูที่ตัดสินใจว่าข้อมูลเก่าจากสถานะก่อนหน้าควรถูกลบหรือยังคงถูกใช้ในการคำนวณสถานะใหม่ โดยกำหนดน้ำหนักว่าข้อมูลส่วนใดควรถูกลบออกหรือเก็บไว้สำหรับการใช้งานต่อไป



ภาพประกอบ 2 โครงสร้างภายในของ Gated Recurrent Unit (GRU)

ที่มา : (Wikipedia, 2024b)

คุณสมบัติเด่นของ GRU

โครงสร้างที่เรียบง่ายกว่า LSTM: ใช้ gate เพียง 2 ตัว (Update Gate และ Reset Gate) แทนที่จะเป็น 3 ตัวเหมือน LSTM (Forget Gate, Input Gate และ Output Gate) ซึ่งช่วยลดจำนวนพารามิเตอร์ ทำให้คำนวณได้เร็วขึ้น

ลดปัญหา Vanishing Gradient: GRU ช่วยรักษาการไหลของข้อมูลย้อนหลังในลำดับข้อมูลยาว ทำให้เรียนรู้ความสัมพันธ์ระยะยาวได้ดีขึ้น

ประสิทธิภาพสูง: ในงานวิจัย (Fu et al., 2016) (Dey & Salem, 2017) GRU แสดงประสิทธิภาพที่ใกล้เคียงกับ LSTM หรือดีกว่าในบางกรณีโดยเฉพาะงานที่ต้องการการคำนวณที่รวดเร็ว เช่น การพยากรณ์ลำดับข้อมูล (Sequence prediction)

2.6 ทฤษฎีเกี่ยวกับ BiLSTM (Bidirectional Long Short-Term Memory)

Bidirectional Long Short-Term Memory (BiLSTM) เป็นการต่อยอดจากโครงข่าย LSTM (Long Short-Term Memory) ซึ่งเป็นรูปแบบหนึ่งของโครงข่ายประสาทเทียมเชิงลำดับ (Recurrent Neural Network: RNN) ที่มีความสามารถในการจดจำข้อมูลลำดับยาวและเข้าใจบริบทในระยะยาวได้ดีกว่า RNN แบบดั้งเดิม (Devlin et al., 2019)

ในขณะที่ LSTM ทำการประมวลผลลำดับข้อมูลจากซ้ายไปขวา (forward direction) เพียงทิศทางเดียว BiLSTM จะมีโครงข่าย LSTM สองชุด ที่ทำงาน ในทิศทางตรงกันข้าม ได้แก่

1. Forward LSTM: ประมวลผลจากคำแรกไปคำสุดท้าย
2. Backward LSTM: ประมวลผลจากคำสุดท้ายไปคำแรก

ผลลัพธ์จากทั้งสองทิศทางจะถูกนำมารวมกัน ทำให้โมเดลสามารถ เข้าใจบริบททั้งก่อนหน้า และถัดไป ได้อย่างครบถ้วน ซึ่งมีความสำคัญอย่างยิ่งในงานประมวลผลภาษาธรรมชาติ (NLP)

สมมติว่าเรากำลังวิเคราะห์ความรู้สึกจากข้อความ “ห้องพักระมัด แต่พนักงานพูดจาไม่ดี” หากใช้ LSTM แบบทิศทางเดียว (เช่น จากซ้ายไปขวา) โมเดลอาจเข้าใจข้อความว่าเป็นเชิงบวกจากคำว่า “สะอาด” ก่อนที่จะอ่านถึงคำว่า “แต่” แต่ใน BiLSTM โมเดลจะอ่านจากต้นประโยคไปท้ายประโยค (เพื่อดูว่าผู้พูดเริ่มจากอะไร) และ อ่านจากท้ายประโยคไปต้นประโยค (เพื่อดูว่าผู้พูดลงท้ายอย่างไร) ทำให้เข้าใจได้ว่า “คำว่า ‘แต่’ เปลี่ยนน้ำเสียงของประโยค” จากบวกเป็นลบ ดังนั้น BiLSTM จึงสามารถ เข้าใจความหมายโดยรวมของประโยคได้แม่นยำกว่า

จุดเด่นของ BiLSTM

1. เข้าใจบริบทใน สองทิศทาง: ทั้งอดีต (past context) และอนาคต (future context)
2. เหมาะสำหรับภาษาไทยที่มีโครงสร้างประโยคยืดหยุ่นและไม่มีการเว้นวรรค
3. ใช้ได้ดีกับงานลำดับ เช่น Named Entity Recognition (NER), Sentiment Analysis, Part-of-Speech Tagging, และ Machine Translation

2.7 ทฤษฎีเกี่ยวกับ WangchanBERTa

WangchanBERTa เป็นโมเดลภาษาไทยเชิงลึกที่พัฒนาขึ้นโดย AI Research Center แห่งสถาบันวิจัยปัญญาประดิษฐ์แห่งชาติ (VISTEC) ซึ่งมีเป้าหมายเพื่อรองรับการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) ในภาษาไทยโดยเฉพาะ เนื่องจากโมเดลภาษาอย่าง BERT แบบดั้งเดิมที่พัฒนาโดย Google นั้นได้รับการฝึกจากข้อมูลภาษาอังกฤษ จึงไม่สามารถประมวลผลภาษาไทยได้อย่างมีประสิทธิภาพเท่าที่ควร (Lowphansirikul et al., 2021)

WangchanBERTa ถูกออกแบบให้มีพื้นฐานจากสถาปัตยกรรม RoBERTa (Robustly Optimized BERT Pretraining Approach) (Liu et al., 2019) ซึ่งเป็นเวอร์ชันปรับปรุงของ BERT โดยมีคุณสมบัติเด่นดังนี้

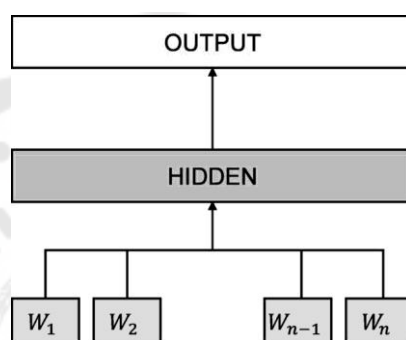
1. รองรับภาษาไทยโดยตรง ถูกฝึกจากข้อมูลภาษาไทยจำนวนมหาศาล (มากกว่า 16,000 ล้าน subwords) จากหลากหลายโดเมน เช่น ข่าวสาร โซเชียลมีเดีย และเว็บไซต์
 2. ใช้ Tokenizer แบบ SentencePiece ซึ่งเหมาะกับภาษาไทยที่ไม่มีการเว้นวรรคระหว่างคำ
 3. มีหลายขนาดโมเดล ให้เลือกใช้งาน เช่น WangchanBERTa-base, WangchanBERTa-large และเวอร์ชันที่ fine-tune สำหรับงานเฉพาะ เช่น sentiment analysis หรือ question answering
- ด้วยความสามารถในการเข้าใจบริบทเชิงลึกของภาษาไทย WangchanBERTa จึงถูกนำมาใช้ในงานวิเคราะห์ข้อความที่ซับซ้อน เช่น การวิเคราะห์ความรู้สึก (Sentiment Analysis) การสรุปข้อความ (Summarization) และการตอบคำถาม (Question Answering)

WangchanBERTa ถือเป็นความก้าวหน้าสำคัญในการพัฒนาโมเดล NLP สำหรับภาษาไทย และเป็นเครื่องมือสำคัญสำหรับการประยุกต์ใช้เทคโนโลยีปัญญาประดิษฐ์กับข้อมูลภาษาไทยอย่างแม่นยำและลึกซึ้ง

2.8 ทฤษฎีเกี่ยวกับ FastText

FastText เป็นเทคนิคการฝังคำ (Word Embedding) และการจัดประเภทข้อความ (Text Classification) ที่พัฒนาโดยทีมวิจัยของ Facebook AI Research (FAIR) ในปี ค.ศ. 2016 โดย Joulin, Grave, Bojanowski, และ Mikolov (Bojanowski et al., 2017)

จุดประสงค์หลักของการพัฒนา FastText คือเพื่อสร้างโมเดลที่สามารถเรียนรู้ความหมายของคำได้อย่างรวดเร็ว มีประสิทธิภาพ และสามารถนำไปใช้งานกับข้อมูลขนาดใหญ่หรือหลายภาษาได้ดี โดยเฉพาะในบริบทของภาษาเชิงวิเคราะห์ เช่น ภาษาไทย ซึ่งไม่มีการเว้นวรรคระหว่างคำ



ภาพประกอบ 3 สถาปัตยกรรมของ FastText ที่ใช้ในงาน text classification

ที่มา : (Wu et al., 2020)

สถาปัตยกรรมของ FastText

ภาพประกอบ 3 เป็นสถาปัตยกรรมพื้นฐานของโมเดล FastText สำหรับการจำแนกประโยคหนึ่งประโยค โดยประกอบด้วยองค์ประกอบสำคัญดังนี้

1. Input ($W_1, W_2, \dots, W_n$)

ข้อความที่นำเข้า (input sentence) จะถูกแบ่งเป็นหน่วยย่อยที่เรียกว่า n-gram features ซึ่งอาจเป็นคำ (word) หรือ subword (เช่น กลุ่มอักขระ เช่น lo, ove จากคำว่า "love")

การใช้ n-gram ช่วยให้ FastText สามารถเข้าใจคำที่ไม่เคยพบมาก่อนได้ดียิ่งขึ้น (out-of-vocabulary words)

2. Embedding Layer $\rightarrow$ Hidden Layer

แต่ละ n-gram จะถูกแปลงเป็นเวกเตอร์ตัวเลขผ่าน embedding matrix จากนั้นเวกเตอร์ของทุก n-gram ในประโยคจะถูกนำมาเฉลี่ย (average) รวมกัน กลายเป็นเวกเตอร์เดี่ยวที่เรียกว่า hidden vector ซึ่งเป็นตัวแทนของทั้งประโยค

3. Output Layer

เวกเตอร์ hidden จะถูกป้อนเข้าไปยังชั้น Output ซึ่งเป็น layer แบบ linear (ไม่มี activation function) แล้วจึงใช้ softmax (หรือ sigmoid ในกรณี multi-label) เพื่อคำนวณค่าความน่าจะเป็นของคลาสต่างๆ เช่น ทำนายว่ารีวิวนั้นเป็น positive / negative / neutral

FastText พัฒนาต่อยอดจากโมเดล Word2Vec โดยเพิ่มคุณสมบัติสำคัญที่เรียกว่า Subword Information ซึ่งช่วยให้สามารถฝึกเวกเตอร์ของคำที่ไม่เคยปรากฏในชุดข้อมูลฝึก (Out-of-Vocabulary words) ได้จากส่วนประกอบของคำนั้น ๆ โดยกลไกการทำงานของ FastText แตกต่างจาก Word2Vec ตรงที่แทนที่จะฝึกเวกเตอร์ของทั้งคำเพียงคำเดียว (เช่น “สะอาด”) โมเดล FastText จะแตกคำออกเป็น subword (n-grams) คือ “สะ” “สะอ” “ะอา” “อาด” “าด” เมื่อโมเดลต้องเรียนรู้ความหมายของคำ “สะอาด” จะเรียนรู้จาก subwords ทั้งหมดเหล่านี้รวมกัน ทำให้สามารถฝึกโมเดลที่เข้าใจคำที่สะกดคล้ายกัน หรือมีโครงสร้างภายในใกล้เคียงกันได้ดีขึ้น ทั้งนี้ ประโยชน์ของการใช้ subword เช่น ช่วยให้เข้าใจความหมายของคำที่ไม่เคยเจอในชุดข้อมูล รองรับภาษาไทยที่มีลักษณะผสมคำได้ดี รวมถึงเพิ่มความยืดหยุ่นในการทำ embedding สำหรับภาษาอื่นๆ

ตัวอย่างการใช้งาน FastText

“ห้องพักระมัด ทำเลดี ใกล้ชายหาด แต่พนักงานพูดจาไม่สุภาพ”

เมื่อใช้ FastText ฝึกเวกเตอร์จากรีวิวเหล่านี้ คำอย่าง “สะอาด” และ “ทำเลดี” จะมีเวกเตอร์ที่อยู่ใกล้กับคำอื่นที่มีความหมายคล้ายกัน เช่น “สะอาดมาก” “ปลอดภัย” “สงบ” ในขณะที่คำอย่าง “ไม่สุภาพ” จะถูกฝึกให้แยกออกไปคนละทิศทางในเวกเตอร์สเปซ ซึ่งช่วยให้โมเดล เรียนรู้ความรู้สึกเชิงบวกและลบได้จากโครงสร้างคำย่อย และสามารถวิเคราะห์ความคิดเห็นได้แม่นยำยิ่งขึ้น แม้จะมีการใช้คำหลากหลายรูปแบบ

ข้อจำกัดของ FastText เช่น

1. จะไม่สามารถเรียนรู้บริบทเชิงลำดับของคำได้ เช่น ประโยคที่มีการเปลี่ยนความหมายด้วยคำว่า “แต่”

2. ไม่สามารถจับความสัมพันธ์ของคำในประโยคได้ดีเท่าโมเดล contextual เช่น BERT หรือ WangchanBERTa

อย่างไรก็ตาม FastText ยังคงเป็นทางเลือกที่ดีมากสำหรับงานที่ต้องการความเร็วและแม่นยำในระดับคำ เช่น การจัดประเภทข้อความ (text classification) การหาคำคล้าย (similarity) หรือ sentiment analysis เบื้องต้น

2.9 ทฤษฎีเกี่ยวกับการประเมินวัดความแม่นยำของแต่ละอัลกอริทึม

ในการพัฒนาอัลกอริทึมการเรียนรู้ของเครื่อง การวัดความแม่นยำ (Accuracy) และประสิทธิภาพของโมเดลเป็นขั้นตอนที่สำคัญในการประเมินว่าโมเดลที่สร้างขึ้นสามารถทำงานได้ดีแค่ไหนภายใต้เงื่อนไขที่กำหนด โดยการวัดความแม่นยำของอัลกอริทึมจะช่วยให้สามารถเปรียบเทียบประสิทธิภาพของโมเดลต่าง ๆ เพื่อเลือกโมเดลที่มีความแม่นยำสูงสุดในการใช้งานที่ต้องการ

Metrics หลักในการวัดความแม่นยำ โดยที่

- True Positive (TP) คือ จำนวนที่ทำนายว่าเป็น Positive และถูกต้อง
- False Positive (FP) คือ จำนวนที่ทำนายว่าเป็น Positive แต่ผิด
- True Negative (TN) คือ จำนวนที่ทำนายว่าเป็น Negative และถูกต้อง
- False Negative (FN) คือ จำนวนที่ทำนายว่าเป็น Negative แต่ผิด

1. Accuracy (ความแม่นยำ) วัดความถูกต้องของโมเดลโดยการคำนวณจากจำนวนการทำนายที่ถูกต้องหารด้วยจำนวนทั้งหมดของการทำนาย ใช้ได้ดีในกรณีที่ข้อมูลในแต่ละคลาสมีสัดส่วนที่ใกล้เคียงกัน ดังสมการ (10)

$$Accuracy = \frac{TN + TP}{TP + FP + TN + FN} \quad (10)$$

2. Precision (ความแม่นยำเฉพาะกลุ่ม) วัดความแม่นยำของโมเดลในกลุ่มข้อมูลที่โมเดลคาดการณ์ว่าเป็นกลุ่มเป้าหมาย (Positive) ใช้ได้ดีในกรณีที่ต้องการลดจำนวนการทำนายผิดที่ว่าเป็น Positive แต่จริง ๆ แล้วเป็น Negative ดังสมการ (11)

$$Precision = \frac{TP}{TP + FP} \quad (11)$$

3. Recall (ค่าความสามารถในการดึงข้อมูลกลับ) วัดความสามารถของโมเดลในการค้นหาตัวอย่างที่เป็นกลุ่มเป้าหมายทั้งหมด ใช้ในกรณีที่ต้องการจับทุกตัวอย่างที่มีความสำคัญให้ครบถ้วน ดังสมการ (12)

$$Recall = \frac{TP}{TP + FN} \quad (12)$$

4. F1 Score เป็นค่าเฉลี่ยแบบถ่วงน้ำหนักของ Precision และ Recall ซึ่งช่วยให้สามารถพิจารณาความสมดุลระหว่าง Precision และ Recall โดยเฉพาะในกรณีที่ข้อมูลไม่สมดุล ดังสมการ (13)

$$F1\ Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (13)$$

2.10 งานวิจัยที่เกี่ยวข้อง

เพื่อให้การดำเนินงานวิจัยในครั้งนี้มีความรอบด้านและตั้งอยู่บนพื้นฐานขององค์ความรู้ที่เกี่ยวข้อง จึงได้ศึกษางานวิจัยที่เกี่ยวข้องทั้งในและต่างประเทศ ซึ่งมีจุดมุ่งหมายคล้ายคลึงกันในด้าน การวิเคราะห์ความรู้สึกจากข้อความ โดยเฉพาะในบริบทของภาษาไทยหรือภาคธุรกิจที่มีลักษณะ คล้ายคลึงกัน เช่น บทความรีวิวสินค้า รีวิวโรงแรม และความคิดเห็นในโซเชียลมีเดีย งานวิจัยเหล่านี้ ครอบคลุมการประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่อง และการเรียนรู้เชิงลึก รวมถึงการใช้เทคนิคการ แปลงข้อความเป็นเวกเตอร์ เช่น TF-IDF, Word2Vec, FastText และโมเดลแบบ Transformer เพื่อเพิ่ม ความแม่นยำในการจำแนกอารมณ์เชิงบวก เชิงลบ หรือเป็นกลางจากข้อความประเภทต่าง ๆ นอกจากนี้ ยังมีการนำอัลกอริทึมที่หลากหลายมาเปรียบเทียบประสิทธิภาพ เช่น Logistic Regression, Support Vector Machine (SVM), Long Short-Term Memory (LSTM), Bidirectional LSTM (BiLSTM), และ ELECTRA เพื่อตรวจสอบความเหมาะสมของแต่ละเทคนิคต่อบริบทของข้อมูลที่แตกต่างกันทั้งในด้าน ภาษา แพลตฟอร์ม และประเภทของข้อความ ทั้งนี้ งานวิจัยที่ได้รับการรวบรวมและนำเสนอในหัวข้อ ต่อไปนี้ จะเป็นประโยชน์ในการวางแผนทางการพัฒนาแบบจำลองของงานวิจัยฉบับนี้ รวมถึงเป็นข้อมูล เปรียบเทียบผลลัพธ์ของโมเดลที่ถูกออกแบบ

2.10.1 งานวิจัยเรื่องตัวแบบการวิเคราะห์ความรู้สึกทางอารมณ์สำหรับจำแนก ประเภทบทความแนะนำสินค้าออนไลน์ (Bowornlertsutee & Paireekreng, 2022)

งานวิจัยเรื่องตัวแบบการวิเคราะห์ความรู้สึกทางอารมณ์สำหรับจำแนกประเภทบทความ แนะนำสินค้าออนไลน์ มีเป้าหมายเพื่อสร้างตัวแบบการวิเคราะห์ความรู้สึกทางอารมณ์จากบทความ แนะนำสินค้าออนไลน์ที่เขียนด้วยภาษาไทย โดยจำแนกเป็น 3 ระดับ ได้แก่ เชิงบวก เป็นกลาง และเชิง

ลบ รวมถึงเปรียบเทียบประสิทธิภาพของเทคนิคการจำแนกประเภทด้วยวิธีการเรียนรู้ของเครื่อง (Machine Learning) แบบมีผู้สอน (Supervised Learning) โดยใช้ชุดข้อมูลความคิดเห็นภาษาไทย จำนวน 12,900 รายการ จากคลังข้อมูล Wisersight Sentiment Corpus ซึ่งมีการติดป้ายกำกับความรู้สึกไว้เรียบร้อยแล้ว ข้อมูลถูกสุ่มคัดเลือกให้มีจำนวนเท่ากันในแต่ละระดับของความรู้สึก โดยการประมวลผลประกอบด้วย 5 ขั้นตอนหลัก ได้แก่

1. เตรียมข้อมูลและทำความสะอาดข้อความ เช่น ตัด HTML URL หรือเครื่องหมายพิเศษ
2. ตัดคำด้วย PyThaiNLP โดยใช้ Attacut Engine และลบคำที่ไม่มีมีความหมาย (stop words)
3. แปลงข้อความเป็นเวกเตอร์ด้วยเทคนิค Bag of Words
4. ฝึกอบรมและทดสอบด้วยโมเดล Machine Learning จำนวน 4 แบบ ได้แก่ LSTM, Logistic Regression, SGD และ SVM
5. ประเมินประสิทธิภาพด้วยค่า Accuracy, Precision, Recall และ F1 Score

จากการทดลองดังกล่าวพบว่า โมเดลที่ใช้เทคนิค Long Short-Term Memory (LSTM) ให้ผลลัพธ์ที่ดีที่สุด โดยมีความถูกต้องในการจำแนกประเภทสูงถึง 81.27% รองลงมา คือ Logistic Regression (69%), Stochastic Gradient Descent หรือ SGD (66%) และ SVM (65%) ตามลำดับ ผลการศึกษาชี้ว่าการใช้ LSTM ซึ่งเป็นเทคนิค Deep Learning บนพื้นฐานของโครงข่ายประสาทเทียม ให้ความสามารถที่เหนือกว่าการจำแนกข้อความภาษาไทยทั่วไป เนื่องจากสามารถเข้าใจบริบทที่ซับซ้อนและโครงสร้างภาษาได้ดีกว่า อย่างไรก็ตาม ควรมีการพัฒนาตัวแบบให้สามารถเข้าใจคำประชดประชัน (sarcasm) และคำแสลงในภาษาไทยได้ดีขึ้นในอนาคต รวมถึงอาจเสริมด้วยเทคนิคการวิเคราะห์ระดับคำ (sentiment word level) และ named-entity tagging เพื่อเพิ่มความแม่นยำของการจำแนกประเภท นอกจากนี้ ตัวแบบที่พัฒนาในงานวิจัยข้างต้นยังสามารถขยายไปใช้กับข้อมูลความคิดเห็นในบริบทอื่น ๆ เช่น โรงแรม ร้านอาหาร หรือบริการออนไลน์ เพื่อสนับสนุนการตัดสินใจของผู้บริโภคและการพัฒนาสินค้าของผู้ประกอบการได้อย่างมีประสิทธิภาพ

2.10.2 งานวิจัยเรื่อง Twitter data sentiment analysis of tourism in Thailand during the COVID-19 pandemic using machine learning (Leelawat et al., 2022)

งานวิจัยเรื่อง Twitter data sentiment analysis of tourism in Thailand during the COVID-19 pandemic using machine learning มีเป้าหมายเพื่อวิเคราะห์ความคิดเห็นของนักท่องเที่ยวชาวต่างชาติที่กล่าวถึงประเทศไทยผ่านแพลตฟอร์ม Twitter ในช่วงเดือนกรกฎาคมถึงธันวาคม พ.ศ. 2563 โดยใช้เทคนิคการเรียนรู้ของเครื่อง (Machine Learning) ในการจำแนก “ความรู้สึก” (Sentiment)

และ “เจตนา” (Intention) ของผู้เขียนทวิต เพื่อนำข้อมูลไปใช้เป็นแนวทางในการฟื้นฟูภาคการท่องเที่ยวของประเทศ โดยกระบวนการศึกษาประกอบด้วย การเก็บข้อมูลทวิตภาษาอังกฤษที่กล่าวถึงสถานที่ท่องเที่ยวสำคัญของไทย ได้แก่ กรุงเทพมหานคร เชียงใหม่ และภูเก็ต จากนั้นคัดกรองข้อมูลที่ไม่เกี่ยวข้อง และประมวลผลด้วยเทคนิคการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) โดยใช้ TF-IDF และการ lemmatization ก่อนนำข้อมูลเข้าสู่การฝึกโมเดล Machine Learning ได้แก่ Support Vector Machine (SVM), Decision Tree (CART) และ Random Forest เพื่อเปรียบเทียบประสิทธิภาพในการทำนาย

ผลการศึกษาข้างต้นแสดงให้เห็นว่า SVM ให้ผลแม่นยำสูงที่สุดในการวิเคราะห์ความรู้สึก โดยเฉพาะกรณีของภูเก็ตซึ่งมีความแม่นยำสูงถึง 77.4% ขณะที่การวิเคราะห์เจตนาในการมาเที่ยวพบว่า Random Forest มีประสิทธิภาพสูงสุดในกรุงเทพฯ ด้วยความแม่นยำ 95.4% นอกจากนี้ การวิเคราะห์เนื้อหา (Content Analysis) ยังพบว่า ความรู้สึกเชิงบวกของนักท่องเที่ยวมักเกี่ยวข้องกับอาหาร สถานที่ท่องเที่ยวทางธรรมชาติ และความสะอาดสวยงาม ในขณะที่ความรู้สึกเชิงลบมักสะท้อนถึงความกังวลต่อสถานการณ์การเมืองและการแพร่ระบาดของ COVID-19 ซึ่งถือเป็นอุปสรรคสำคัญต่อการตัดสินใจเดินทางมายังประเทศไทย อย่างไรก็ตาม หากต้องการฟื้นฟูการท่องเที่ยวไทยอย่างยั่งยืน ควรให้ความสำคัญกับการลดความกังวลของนักท่องเที่ยวเกี่ยวกับ COVID-19 และปัญหาทางการเมือง รวมถึงส่งเสริมจุดแข็งของประเทศใน 4 ด้าน ได้แก่ สถานที่ท่องเที่ยว ธรรมชาติ อาหาร และกิจกรรมยามค่ำคืน นอกจากนี้ หน่วยงานที่เกี่ยวข้องควรพิจารณาใช้ข้อมูลจากโซเชียลมีเดียเป็นเครื่องมือในการวิเคราะห์ความรู้สึกและความต้องการของนักท่องเที่ยว เพื่อกำหนดนโยบายและกลยุทธ์การตลาดได้อย่างตรงเป้าหมาย

2.10.3 งานวิจัยเรื่อง Sentiment Analysis for Thai Language in Hotel Domain Using Machine Learning Algorithms (Khamphakdee & Seresangtakul, 2021)

งานวิจัยเรื่อง Sentiment Analysis for Thai Language in Hotel Domain Using Machine Learning Algorithms มีเป้าหมายเพื่อพัฒนาระบบวิเคราะห์ความรู้สึก (Sentiment Analysis) สำหรับภาษาไทยในบริบทของรีวิวโรงแรม โดยใช้เทคนิคการเรียนรู้ของเครื่อง (Machine Learning: ML) เพื่อจำแนกรีวิวของลูกค้าให้เป็นความคิดเห็นเชิงบวกหรือเชิงลบอย่างแม่นยำ เหมาะสมสำหรับนำไปประยุกต์ใช้ในธุรกิจโรงแรมขนาดเล็กและขนาดกลาง (SMEs) ในประเทศไทย ซึ่งข้อมูลที่ใช้ในการศึกษาถูกรวบรวมจากเว็บไซต์ Agoda.com และ Booking.com รวมเป็นจำนวน 16,804 รีวิว จากลูกค้าที่ใช้ภาษาไทย และผ่านกระบวนการคัดเลือกรีวิวจำนวน 1,000 รายการ เพื่อให้ผู้เชี่ยวชาญแปลป้ายกำกับ

ความรู้สึก (positive/negative) สำหรับสร้างชุดข้อมูลฝึก (training corpus) และใช้เทคนิค cosine similarity ในการสร้าง corpus สำหรับรีวิวที่เหลือ โดยขั้นตอนการประมวลผลประกอบด้วย

1. การทำความสะอาดข้อมูล (data preprocessing) รวม 11 ขั้นตอน เช่น ตัดคำ ตรวจสอบคำที่ผิด และตัดอีโมจิ ฯลฯ
2. การแปลงข้อความเป็นเวกเตอร์ด้วยเทคนิค TF-IDF, Delta TF-IDF, N-Gram (unigram, bigram, trigram) และ Word2Vec
3. การจำแนกความคิดเห็นโดยใช้ ML ทั้งหมด 9 แบบ เช่น SVM, Random Forest, Logistic Regression, Naïve Bayes ฯลฯ พร้อมประเมินด้วย k-fold cross-validation (k=10)

จากผลการศึกษาดังกล่าว พบว่า เทคนิคที่ได้ผลดีที่สุดคือ Support Vector Machine (SVM) ร่วมกับ Delta TF-IDF ซึ่งให้ความแม่นยำสูงสุดถึง 89.96% เช่นเดียวกับเทคนิค Word2Vec (skip-gram) ร่วมกับ Ridge Regression ให้ผลแม่นยำดีเช่นกันที่ 87.82% ในขณะที่ N-Gram แบบ trigram ให้ผลลัพธ์ต่ำที่สุดเมื่อเทียบกับเทคนิคอื่น โดยเฉพาะเมื่อขนาดข้อมูลฝึกมีจำนวนน้อย และเมื่อเปรียบเทียบแบบจำลองทั้ง 9 แบบ แสดงให้เห็นว่า การเลือกเทคนิคแปลงข้อความที่เหมาะสมมีผลต่อความแม่นยำของโมเดลอย่างชัดเจน รวมทั้งการวิเคราะห์ความรู้สึกในภาษาไทยจำเป็นต้องมีการเตรียมข้อมูลที่รอบคอบ โดยเฉพาะขั้นตอนการตัดคำและแก้คำที่ผิด ซึ่งมีผลโดยตรงต่อผลลัพธ์ของโมเดล ดังนั้น จึงควรมีการพัฒนา word embedding แบบ pre-trained สำหรับภาษาไทยที่ครอบคลุมมากขึ้น เช่น Thai2Vec, BERT ฯลฯ รวมถึงการใช้ deep learning ในอนาคต เพื่อยกระดับประสิทธิภาพของระบบวิเคราะห์ความรู้สึกในภาคการบริการและท่องเที่ยว

2.10.4 งานวิจัยเรื่อง Sentiment Analysis of Hotel Reviews With LSTM And ELECTRA (Husein et al., 2023)

งานวิจัยเรื่อง Sentiment Analysis of Hotel Reviews With LSTM And ELECTRA มีจุดมุ่งหมายเพื่อพัฒนาโมเดลวิเคราะห์ความรู้สึกจากข้อความรีวิวโรงแรม โดยเน้นการเปรียบเทียบประสิทธิภาพของอัลกอริทึม LSTM (Long Short-Term Memory) และ ELECTRA (Efficiently Learning an Encoder that Classifies Token Replacements Accurately) ในการจำแนกรีวิวเป็นความรู้สึกเชิงบวก เชิงลบ หรือเป็นกลาง ข้อมูลที่นำมาใช้เป็นรีวิวจากโรงแรม Grand Mercure Maha Cipta Medan Angkasa ในประเทศอินโดนีเซีย ซึ่งถือเป็นชุดข้อมูลใหม่ที่ยังไม่เคยถูกศึกษามาก่อน โดยใช้ข้อมูลรีวิวจำนวน 977 รายการถูกรวบรวมจากเว็บไซต์ TripAdvisor ผ่านเครื่องมือ Webscraper.io ข้อมูลได้ถูกประมวลผลเบื้องต้นด้วยกระบวนการล้างข้อความ เช่น การลบเครื่องหมายวรรคตอน เปลี่ยนตัวอักษรให้

เป็นพิมพ์เล็ก ลบ stopwords และแปลงคำให้อยู่ในรูปคำพื้นฐาน (lemmatization) หลังจากนั้นจัดสมดุลข้อมูลด้วยเทคนิค undersampling เพื่อให้จำนวนรีวิวในแต่ละประเภทของความรู้สึกใกล้เคียงกัน ก่อนนำเข้าสู่กระบวนการแบ่งข้อมูล (Train-Validation-Test) และวิเคราะห์ด้วยอัลกอริทึม LSTM และ ELECTRA ตามลำดับ

จากการเปรียบเทียบประสิทธิภาพของโมเดลทั้งสอง พบว่า ELECTRA ให้ความแม่นยำสูงกว่า โดยมีค่า Accuracy ที่ 47% และ F1-score อยู่ในช่วง 0.35 - 0.59 และ LSTM มีค่า Accuracy เพียง 30% และ F1-score อยู่ในช่วง 0.28 - 0.33 ทั้งนี้ ELECTRA ยังมีข้อได้เปรียบด้านประสิทธิภาพ เพราะไม่จำเป็นต้องใช้ Tokenizer เหมือน LSTM จึงประมวลผลได้เร็วกว่า ซึ่งแสดงให้เห็นว่า ELECTRA มีศักยภาพมากกว่าในการวิเคราะห์ความรู้สึกจากข้อความภาษาอินโดนีเซียในบริบทของรีวิวโรงแรม และสามารถนำไปปรับใช้กับงานที่ต้องการการประมวลผลข้อความขนาดใหญ่ในอนาคต ดังนั้น จึงควรให้ใช้ ELECTRA เป็นพื้นฐานในการพัฒนาโมเดลวิเคราะห์ความรู้สึกในอนาคต โดยเฉพาะเมื่อมีข้อจำกัดด้านเวลาและทรัพยากรการประมวลผล นอกจากนี้ยังสามารถขยายการวิจัยเพื่อทดลองกับข้อมูลจากรีวิวโรงแรมอื่นๆ ในระดับภูมิภาค หรือพัฒนาโมเดลให้รองรับหลายภาษาเพื่อการประยุกต์ใช้งานในระดับนานาชาติ

2.10.5 งานวิจัยเรื่อง Thai Sentiment Analysis via Bidirectional LSTM-CNN Model with Embedding Vectors and Sentic Features (Ayutthaya & Pasupa, 2018)

งานวิจัยเรื่อง Thai Sentiment Analysis via Bidirectional LSTM-CNN Model with Embedding Vectors and Sentic Features มีเป้าหมายเพื่อพัฒนาและปรับปรุงประสิทธิภาพของการวิเคราะห์ความรู้สึกในข้อความภาษาไทย ด้วยการประยุกต์ใช้โมเดล Deep Learning โดยเฉพาะการรวมโมเดลแบบสองทิศทาง (Bidirectional LSTM) กับโครงข่ายประสาทเทียมแบบคอนโวลูชัน (CNN) และเสริมด้วยคุณลักษณะข้อมูลหลากหลาย ได้แก่ Word Embedding (Thai2Vec), Part-of-Speech Embedding และ Sentic Features เพื่อให้สามารถเข้าใจอารมณ์ในข้อความได้อย่างแม่นยำมากขึ้น ข้อมูลที่ใช้ในการทดลองคือประโยคจากนิทานเด็กภาษาไทยจำนวน 1,115 ประโยค โดยมีการกำกับอารมณ์ไว้ 3 ระดับ ได้แก่ เติบโต เป็นกลาง และเชิงบวก โดยได้รับการเห็นพ้องจากผู้เชี่ยวชาญ 3 ท่าน โดยกระบวนการวิเคราะห์ประกอบด้วย

1. การตัดคำด้วย KUCUT และการระบุชนิดคำด้วย Jitar POS tagger
2. การสร้างเวกเตอร์ฝังความหมายด้วย Thai2Vec (300 มิติ) POS Embedding (50 มิติ) และ Sentic Features จาก SenticNet2 (5 มิติ)

3. การรวมข้อมูลเข้ากับโมเดล BiLSTM-CNN ที่ออกแบบให้เรียนรู้ทั้งลำดับคำและโครงสร้างไวยากรณ์ในข้อความ

4. การวัดผลด้วยค่า F1-score โดยแบ่งข้อมูลเป็นชุด train validation และ test ตามอัตราส่วน 80:10:10

ผลลัพธ์แสดงให้เห็นว่า การใช้เพียง Word Embedding ให้ค่า F1-score ที่ 75.75% และ การใช้ Sentic Feature เพียงอย่างเดียว ให้ค่า F1-score ที่ 69.26% ในขณะที่เมื่อรวม Word Embedding กับ POS Embedding และ Sentic Feature เข้าด้วยกัน โมเดล BiLSTM-CNN สามารถให้ค่าความแม่นยำสูงสุดที่ 78.89% นอกจากนี้ยังพบว่า การเสริมข้อมูลเชิงไวยากรณ์ (POS) และข้อมูลเชิงอารมณ์ (Sentic) ช่วยให้โมเดลเข้าใจความหมายและอารมณ์ในข้อความได้ดีมากยิ่งขึ้น ดังนั้น การพัฒนาแบบจำลองการวิเคราะห์ความรู้สึกในภาษาไทยควรคำนึงถึงมิติด้านความหมาย ไวยากรณ์ และอารมณ์ร่วมกัน โดยเฉพาะในบริบทของภาษาไทยที่ไม่มีการเว้นวรรคระหว่างคำ การรวมคุณลักษณะหลายด้านช่วยเพิ่มประสิทธิภาพได้อย่างชัดเจน

2.10.6 งานวิจัยเรื่อง Sentiment Analysis on TikTok App using Long Short-Term Memory (LSTM) with Stochastic Gradient Descent (SGD) Optimization (Faqia et al., 2024)

งานวิจัยเรื่อง Sentiment Analysis on TikTok App using Long Short-Term Memory (LSTM) with Stochastic Gradient Descent (SGD) Optimization มีเป้าหมายเพื่อวิเคราะห์ความรู้สึกของผู้ใช้งานแอปพลิเคชัน TikTok จากรีวิวบน Google Play Store โดยใช้เทคนิคการเรียนรู้ของเครื่องแบบ Long Short-Term Memory (LSTM) ร่วมกับการปรับค่าพารามิเตอร์ด้วยอัลกอริทึม Stochastic Gradient Descent (SGD) และเทคนิคเพิ่มคุณภาพของข้อมูลด้วย Word2Vec และ FastText เพื่อเพิ่มประสิทธิภาพของโมเดลในการจำแนกความรู้สึกเชิงบวกและเชิงลบอย่างแม่นยำ โดยใช้ข้อมูลรีวิวแอป TikTok จำนวน 15,049 รายการ ซึ่งถูกดึงจาก Google Play Store และผ่านการติดป้ายกำกับเป็น รีวิวเชิงบวก (1) และเชิงลบ (0) โดยผู้เชี่ยวชาญสองท่าน ทั้งนี้ ในการทดลองแบ่งออกเป็น 4 รูปแบบ (Scenarios) ได้แก่:

Scenario 1: ใช้ LSTM พื้นฐาน

Scenario 2: เพิ่ม Word2Vec เพื่อแปลงคำเป็นเวกเตอร์

Scenario 3: เพิ่ม FastText เพื่อเสริมบริบทคำย่อยและคำสะกดผิด

Scenario 4: เพิ่มอัลกอริทึม SGD เพื่อปรับค่าพารามิเตอร์ให้เหมาะสมที่สุดการประเมินผลใช้ ค่า Accuracy, Precision, Recall และ F1-Score ผ่านการทดสอบแบบ Confusion Matrix และแบ่งชุดข้อมูลเป็น Train/Test ด้วยอัตราส่วน 90:10

จากผลการทดลองดังกล่าวพบว่า

Scenario 1 (LSTM): ความแม่นยำ (Accuracy) อยู่ที่ 80.73%

Scenario 2 (LSTM + Word2Vec): เพิ่มขึ้นเป็น 85.91%

Scenario 3 (LSTM + Word2Vec + FastText): ขยับขึ้นเป็น 86.91%

Scenario 4 (LSTM + Word2Vec + FastText + SGD): ให้ผลดีที่สุดที่ Accuracy 88.17% และ F1-Score สูงถึง 88.10%

โดยการเพิ่มการฝังเวกเตอร์และการปรับค่าด้วย SGD ช่วยให้โมเดลเรียนรู้ความสัมพันธ์ของคำได้ดีขึ้น เข้าใจบริบทของข้อความมากขึ้น และจำแนกความรู้สึกได้อย่างแม่นยำมากขึ้น และงานวิจัยดังกล่าวเสนอว่า ควรใช้เทคนิคการฝังความหมายของคำ (word embedding) ร่วมกับอัลกอริทึมการเพิ่มประสิทธิภาพ เช่น SGD เพื่อเพิ่มความแม่นยำในการวิเคราะห์ความรู้สึกจากข้อความ โดยเฉพาะในแอปพลิเคชันที่มีผู้ใช้งานหลากหลายและรีวิวจำนวนมากอย่าง TikTok นอกจากนี้ ควรมีการศึกษาต่อโดยนำโมเดล Deep Learning อื่น ๆ มาเปรียบเทียบ เช่น Bi-LSTM, GRU หรือ Transformer รวมถึงลองใช้เทคนิคการเรียนรู้แบบไม่ต้องมีผู้สอน (unsupervised) กับชุดข้อมูลภาษาอื่นหรือประเภทแพลตฟอร์มอื่น

ภาพรวมของการสรุปงานวิจัยที่เกี่ยวข้อง

ตาราง 5 สรุปภาพรวมงานวิจัยที่เกี่ยวข้อง

ผู้แต่ง	ข้อมูล	ขั้นตอนการ ดำเนินการ	ตัวชี้วัด	ผลลัพธ์
N. Leelawat, S. Jariyapongpai boon, A. Promjun	ทวิต ภาษาอังกฤษ เกี่ยวกับการ ท่องเที่ยวใน กรุงเทพมหานคร เชียงใหม่ และ ภูเก็ต ตั้งแต่ กรกฎาคมถึง ธันวาคม 2020	ใช้อัลกอริทึมการ เรียนรู้ของเครื่อง (Decision Tree, Random Forest, SVM) เพื่อวิเคราะห์ อารมณ์และ ความตั้งใจจาก ทวิต	Accuracy, Precision, Recall, F1- score	การศึกษาพบว่า อารมณ์บวกส่วน ใหญ่เกี่ยวข้อง กับสถานที่ ท่องเที่ยว ธรรมชาติ และ ชีวิตยามค่ำคืน อารมณ์ลบส่วน ใหญ่เกี่ยวกับ COVID-19 และ ความตึงเครียด ทางการเมือง
Nattawat Khamphakdee, Pusadee Seresangtakul	รีวิวจาก Agoda.com และ Booking.com	ใช้ Delta TF- IDF, TF-IDF, N- Gram และ Word2Vec สำหรับการแปลง ข้อความเป็น เวกเตอร์ และเปรียบเทียบ อัลกอริทึม ML ทั้งหมด (9 ตัว) โดยใช้ ten-fold	Recall, precision, F1- score, accuracy	SVM ด้วย Delta TF-IDF มีความ แม่นยำที่สุด ทำ ได้ถึง 89.96%

ผู้แต่ง	ข้อมูล	ขั้นตอนการ ดำเนินการ	ตัวชี้วัด	ผลลัพธ์
		cross-validation		
Pisit Bowornlertsute e, Worapat Paireekreng	12,900 รีวิว ออนไลน์ ภาษาไทย มีป้ายกำกับเป็นบวก กลาง หรือลบ	ใช้เทคนิค Word Segmentation, Bag of Words, และทดสอบกับ 4 โมเดลได้แก่ LSTM, SGD, Logistic Regression และ SVM	accuracy, precision, recall, and F1-score	LSTM มีความแม่นยำสูงสุด 81.27% ซึ่งมีประสิทธิภาพในการจำแนกอารมณ์ได้ดีกว่าโมเดลอื่น
Amir Mahmud Husein, Nicholas Livando, Andika, William Chandra, Gary Phan	รีวิวโรงแรม 977 รายการจาก Tripadvisor สำหรับโรงแรมในสุมาตราเหนือ, อินโดนีเซีย	ใช้โมเดล LSTM และ ELECTRA สำหรับการทำนายอารมณ์ ประมวลผลข้อมูลเพื่อหาความแม่นยำที่ดีขึ้น	Accuracy, where ELECTRA outperformed LSTM with a higher accuracy rate.	ELECTRA มีประสิทธิภาพดีกว่า LSTM ด้วยความแม่นยำ 47% เทียบกับ 30% ของ LSTM
Thititorn Seneewong Na Ayutthaya, Kitsuchart Pasupa	ข้อมูลรีวิว โรงแรมในไทยที่เก็บรวบรวมจากโซเชียลมีเดีย และแพลตฟอร์มรีวิว เช่น Agoda และ TripAdvisor	ใช้โมเดล LSTM, CNN และ GRU พร้อมด้วยเทคนิคการฝังคำ เช่น Word2Vec และ FastText	F1-score	โมเดล LSTM มีประสิทธิภาพดีกว่า CNN และ GRU โดยมีความแม่นยำถึง 85% ในการจำแนกความคิดเห็นของรีวิวโรงแรมไทย

ผู้แต่ง	ข้อมูล	ขั้นตอนการ ดำเนินการ	ตัวชี้วัด	ผลลัพธ์
Muhammad	15,049 TikTok	ใช้โมเดล LSTM	Accuracy,	LSTM +
Zacky Faqia	app reviews	พร้อมด้วย	Precision,	Word2Vec +
Rizky*, Yuliant	from the	เทคนิคการฝังคำ	Recall และ F1-	FastText +
Sibaroni, Sri	Google Play	เช่น Word2Vec	Score	SGD : ให้ผลดี
Suryani		และ FastText		ที่สุดที่
Prasetyowati		และ ใช้ SGD		Accuracy
		เพื่อปรับค่าพารามิเตอร์		88.17% และ
		มีเตอให้		F1-Score สูงถึง
		เหมาะสมที่สุด		88.10%

จากงานวิจัยที่เกี่ยวข้อง ทำให้ได้เรียนรู้แนวทาง เทคนิค และข้อค้นพบที่มีประโยชน์อย่างยิ่งต่อการออกแบบการวิจัย โดยสามารถสรุปประเด็นสำคัญที่นำมาประยุกต์ใช้ได้ดังนี้

1. **ด้านการคัดเลือกแหล่งข้อมูลและการเตรียมข้อมูล** งานวิจัยหลายฉบับ เน้นการใช้ข้อมูลรีวิวที่เป็นภาษาไทยจากแหล่งข้อมูลที่มีความน่าเชื่อถือ โดยเฉพาะแพลตฟอร์มออนไลน์ที่มีผู้ใช้งานจริงรีวิวบริการหรือสินค้าโดยตรงไปตรงมา พร้อมเกณฑ์คัดกรองเพื่อรักษาคุณภาพของข้อมูล เช่น ความเกี่ยวข้องกับสถานที่จริง ความสมบูรณ์ของข้อความ และการใช้ภาษาไทย

2. **ด้านกระบวนการเตรียมข้อมูล** งานวิจัยส่วนใหญ่ให้ความสำคัญกับขั้นตอนการล้างข้อมูล การตัดคำภาษาไทย การลบคำฟุ่มเฟือย และการแปลงข้อความเป็นเวกเตอร์ เช่น TF-IDF, Word2Vec, FastText และการใช้ Tokenizer จาก BERT ซึ่งงานวิจัยนี้มีการประยุกต์ใช้เทคนิค PyThaiNLP ร่วมกับการฝังข้อความด้วย FastText และ WangchanBERTa เพื่อนำเสนอคุณลักษณะที่เหมาะสมก่อนเข้าสู่แบบจำลองการเรียนรู้

3. **ด้านการเลือกแบบจำลอง** หลายงานวิจัย พบว่าโมเดล LSTM, BiLSTM และ GRU ให้ผลลัพธ์ที่ดีในการจำแนกความรู้สึกจากข้อความภาษาไทย โดยเฉพาะเมื่อผสานกับโมเดลฝังคำที่เหมาะสม ในงานวิจัยนี้ จึงได้ทดลองใช้แบบจำลองทั้ง 3 แบบร่วมกับทั้ง FastText และ WangchanBERTa เพื่อเปรียบเทียบประสิทธิภาพในการวิเคราะห์ความรู้สึกจากข้อความรีวิวที่มีโครงสร้างและบริบทซับซ้อน

4. ด้านการประเมินประสิทธิภาพแบบจำลอง งานวิจัยเกือบทั้งหมดใช้ตัวชี้วัดมาตรฐาน เช่น Accuracy, Precision, Recall และ F1-score ในการเปรียบเทียบผลลัพธ์ของโมเดล ซึ่งผู้วิจัยได้นำแนวทางดังกล่าวมาใช้เพื่อวัดประสิทธิภาพของแบบจำลองในแต่ละรูปแบบอย่างรอบด้าน โดยคำนึงถึงทั้งความถูกต้องและความสามารถในการแยกแยะอารมณ์จากข้อความในสถานการณ์จริง

ความรู้จากงานวิจัยที่เกี่ยวข้องไม่เพียงแต่ช่วยยืนยันแนวทางการดำเนินการของงานวิจัยนี้เท่านั้น แต่ยังมีบทบาทสำคัญในการวางรากฐานเชิงทฤษฎีสำหรับการออกแบบกระบวนการวิจัยอย่างเป็นระบบ ทั้งในด้านการเตรียมข้อมูล การสร้างแบบจำลอง และการประเมินผลลัพธ์ในบริบทของข้อความภาษาไทยจากวีวบนแพลตฟอร์มจริง



บทที่ 3

วิธีการดำเนินงาน

ขั้นตอนการดำเนินงานในวิจัยนี้ได้ถูกออกแบบให้สอดคล้องกับวัตถุประสงค์ของการวิจัย โดยมุ่งเน้นการเลือกแหล่งข้อมูลสาธารณะจากแพลตฟอร์ม Google Maps ที่มีความน่าเชื่อถือ การกำหนดเกณฑ์การคัดเลือกรีวิวที่ชัดเจน การจัดการข้อมูลข้อความด้วยเทคนิค NLP ที่รองรับภาษาไทย ตลอดจนการเปรียบเทียบผลลัพธ์ของแบบจำลอง Deep Learning ที่มีประสิทธิภาพสูง เช่น LSTM, BiLSTM และ GRU ร่วมกับเทคนิคการฝังข้อความแบบ FastText และ WangchanBERTa เพื่อตอบโจทย์การวิเคราะห์ความรู้สึกจากข้อความภาษาไทยในรีวิวโรงแรมบนเกาะภูเก็ตอย่างเป็นรูปธรรม พร้อมทั้งประเมินผลด้วยค่าชี้วัดมาตรฐาน ได้แก่ Accuracy, Precision, Recall และ F1-score เพื่อให้สามารถเปรียบเทียบประสิทธิภาพของแต่ละโมเดลได้อย่างรอบด้าน โดยขั้นตอนดำเนินงานวิจัยที่สำคัญประกอบด้วย 6 ขั้นตอนหลัก ได้แก่ 1) แผนการดำเนินงานวิจัย 2) การคัดเลือกและการเก็บรวบรวมข้อมูล 3) กระบวนการทำงานของแบบจำลอง 4) การเตรียมข้อมูล 5) การสร้างแบบจำลอง และ 6) การประเมินผลแบบจำลอง ซึ่งมีรายละเอียดของแต่ละขั้นตอนสำคัญ ดังนี้

3.1 แผนการดำเนินงานวิจัย

ผู้วิจัยทำการวางแผนดำเนินงานวิจัย โดยมีระยะเวลาดำเนินการทั้งหมด 8 เดือน ระหว่างเดือนสิงหาคม พ.ศ. 2567 ถึงเดือน เมษายน พ.ศ. 2568 ซึ่งมีขั้นตอนดำเนินการ ดังตาราง 6

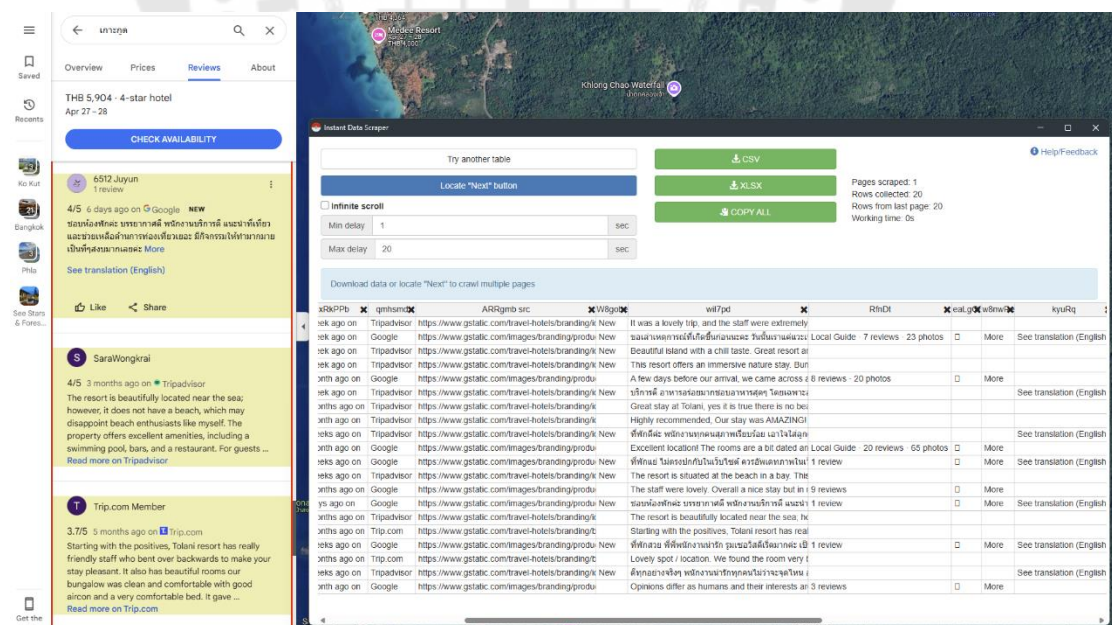
ตาราง 6 แผนการดำเนินงาน

ขั้นตอนการดำเนินงาน	พ.ศ.2567					พ.ศ.2568		
	ส.ค.	ก.ย.	ต.ค.	พ.ย.	ธ.ค.	ม.ค.	ก.พ.	มี.ค.
1. วางแผนการทำงาน								
2. ศึกษาทฤษฎีและแนวคิดที่ใช้								
3. เก็บรวบรวมข้อมูล								
4. เตรียมข้อมูล								
5. สร้างแบบจำลอง								
6. ปรับปรุงประสิทธิภาพแบบจำลอง								
7. ประเมินผลแบบจำลอง								
8. สรุปผลและรายงานผล								

3.2 การเก็บรวบรวมข้อมูลและการคัดเลือก

ในงานวิจัยนี้ ผู้วิจัยได้ดำเนินการเก็บรวบรวมข้อมูลจากแหล่งข้อมูลสาธารณะโดยใช้เทคนิค Web Scraping จากเว็บไซต์ Google Maps (Google, 2025) ซึ่งเป็นแพลตฟอร์มที่ให้บริการแผนที่และรีวิวสถานที่ต่าง ๆ จากผู้ใช้งานจริง โดยมุ่งเน้นข้อมูลจากการรีวิวของโรงแรมและที่พักบนเกาะภูเก็ต จังหวัด ภูเก็ต ซึ่งเป็นหนึ่งในแหล่งท่องเที่ยวยอดนิยมของประเทศไทยที่มีจำนวนนักท่องเที่ยวทั้งชาวไทยและชาวต่างชาติเป็นจำนวนมาก ทำให้สามารถใช้เป็นกรณีศึกษาในการวิเคราะห์ความรู้สึกจากความคิดเห็นของผู้ใช้บริการได้อย่างเหมาะสม

กระบวนการเก็บข้อมูลด้วยเทคนิค Web Scraping ดังกล่าว ดำเนินการโดยใช้เครื่องมืออัตโนมัติซึ่งอยู่ในรูปแบบส่วนขยาย (Extension) ของเว็บเบราว์เซอร์ Google Chrome ที่มีชื่อว่า Instant Data Scraper (ภาพประกอบ 4) โดยเครื่องมือนี้ถูกนำมาใช้ในการเข้าถึงและดึงข้อมูลรีวิวของโรงแรมแต่ละแห่งจากหน้าร้าน (Place Page) ที่ปรากฏอยู่ในระบบของ Google Maps โดยชุดข้อมูลที่ได้ครอบคลุมช่วงเวลา ตั้งแต่วันที่ 1 มกราคม พ.ศ. 2566 ถึงวันที่ 31 มีนาคม พ.ศ. 2567 โดยข้อมูลที่ได้หลังจากทำการ scraping data ยังไม่ผ่านการเตรียมข้อมูลจึงแสดงข้อผิดพลาดที่ไม่สมบูรณ์และยังไม่เหมาะสมต่อการนำไปใช้ (ภาพประกอบ 5)



ภาพประกอบ 4 หน้าจอการใช้งาน Instant Data Scraper

	H	I	J	K	L	M
1	mb src	will7pd	w8nwRe	fontLabelMedium	dSug	Hz
2	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	เป็นรีสอร์ทที่มีทำเลดี ติดหาดที่สวยงาม บรรยากาศดีมาก การบริการตั้งแต่เช็คอินพาเราไปห้อง Reception	เพิ่มเติม	12:10:00 AM		ขอ
3	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	มาที่นี่เป็นครั้งที่ 2 และเป็นครั้งสุดท้ายสำหรับที่นี่แน่นอนเริ่มต้นที่ห้องเรา เจอวิวทะเล หน้าห้อง 1 ตัว	ห้องใหญ่เพิ่มเติม			
4	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพักกว้างสะอาด ทันสมัย / ห้องน้ำกว้าง น่าน้ำ / แต่ในห้องเยอะมากอยู่เยอะ / แอร์เย็นดีมาก	บริเวณไหนเพิ่มเติม			
5	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	มีอาหารเช้าที่อร่อยมาก 7 วันเต็มๆ โดยรวมพนักงานที่เป็นคนกันเอง ดูแลพวกเราและเพื่อน ๆ ได้ดีและ	เพิ่มเติม			
6	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	มาเช้าวันที่ 2 ที่พักดี วิวสวย มีสระว่ายน้ำไม่สวยจริงๆมากก ห้องกว้าง ห้องน้ำกว้าง ราคาไม่แพง	ในรีสอร์ทมีเพิ่มเติม	12:08:00 AM		
7	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องดี ไม่ใหญ่ไม่เล็กเตียงนอนแข็งนิดนึง แต่การดูแลที่ความสะอาดห้องดี ทำให้อยู่ไม่เบื่อจนเกินไป	รับได้ในเพิ่มเติม	12:12:00 AM		
8	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพักดี ห้องสวย เดินไม่ไกลจากทะเลดีตรงมีที่กางผ้าที่ห้อง5ตัว	มีแม่บ้านคอยดูแล และที่นอน ...	เพิ่มเติม		
9	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องสะอาด เงียบสงบ ห้องน้ำใหญ่สะอาด อาหารเช้าของโรงแรมดีมาก ประทับใจ ตอนอาหารเช้าคนน้อย	เพิ่มเติม			
10	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ดีค่ะ ห้องใหญ่สะอาด ส่วนกลางดี ติดหาด แต่ทางเดินไปห้องพักเป็นป่า ตอนกลางคืนแอบอันตราย	เพราะเราเจอเพิ่มเติม			ขอ
11	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพัก : เตียงนุ่ม หมอนดี แต่โดยรวมแล้วถือว่ากลางๆก่อนไปทางพอใช้ แล้วพอได้เราเจอห้อง B3	ที่เหมือนมีเพิ่มเติม			ขอ
12						
13						
14						
15	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพักดี สะอาด ดูไม่เก่า ดูแล้วได้ใจ เครื่องใช้ต่าง ๆ ครบครัน facilities ครบ มีสระว่ายน้ำ มีอุปกรณ์ให้	kayเพิ่มเติม			ขอ
16	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	เป็นรีสอร์ทแบบเกาะภูเก็ตที่มีความเป็นธรรมชาติมาก บ้านพักแยกเป็นหลังคล้ายกระท่อมแต่มีสิ่งอำนวยความสะดวก	เพิ่มเติม			
17	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ทุกอย่างโอเคหมด บรรยากาศดี เงียบสงบวิวสวยวิวดี โรงแรมได้ใจ จะมีข้อสังเกตอยู่ถ้าใครชอบที่เงียบ	เพิ่มเติม	12:04:00 AM		ขอ
18	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพักดีหลายระดับราคา ดีดีหาดแพงสุด ห้องกว้างวิวสวย แดดหน้ามึงคงสวยไม่ลึก ที่นี่ย่อย ห้องน้ำกว้าง	เพิ่มเติม			
19	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	น้ำสะอาดวิวเปลี่ยนที่ที่ข่อย ไม่ควรนำกลับมาใช้อีก ระบบระบายอากาศ ระบบน้ำในห้องพักต้องปรับปรุง	เพิ่มเติม			
20	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ครั้งแรกพอ สำหรับที่นี่ พลาเองที่เชื่อวิวจากในฟีดBACK มีอยู่อย่างเดียววิวทะเลตรงส่วนกลาง แต่ข้างส่วน	เพิ่มเติม			
21	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	เดินทางจากท่าเรือมาสะดวกมาก ทางขึ้นเขาเยอะแต่ไม่ต้องกลัวรถขึ้นเขาจอดอยู่ดีทุกพอ ที่พักดี	เพิ่มเติม			ขอ
22	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพักดีสะอาด ห้องใหญ่สะอาด ได้ห้องวิวมาก ไม่เกินครึ่งครึ่งได้พักผ่อนค่ะ เป็นรีสอร์ทที่มีครบทั้ง	บาร์ คาเฟ่เพิ่มเติม			ขอ
23	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ที่พักมีอาหารเช้าและเครื่องดื่มมากมายที่กินเป็นส่วนตัว แอ่อาหารแบบแพง(ตามราคาของเกาะ)	อาหารเพิ่มเติม			ขอ
24	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	วิวจากหน้าอาหารและวิวรอบๆวิวสวยมาก สามารถลงเล่นน้ำได้ รร. ได้เลย ตอนนั่งเล่นไม่ผิดหวัง	วิวเพิ่มเติม			ขอ
25	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ห้องพักสะอาด กว้างขวางของไป 4 ห้อง เราได้ห้องดีสุด เพื่อนที่เหลือน้อยคน แอร์ไม่เย็น ไม่มีบันไดขึ้นห้อง	เพิ่มเติม	12:09:00 AM		ขอ
26	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	ที่พักดีสะอาด วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย วิวสวย	เพิ่มเติม			
27	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	วิวสวย ตอนไปจองเช็คที่อาหารทะเลดีมาก แต่ที่นี่ย่านมาก ๆ ครั้งจริงไม่ได้สำหรับยาและสติกเกอร์	เพิ่มเติม			ขอ
28	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	Schönes Resort in traumbarer Lage, aber in die Jahre gekommen und dementsprechend heruntergekommen	เพิ่มเติม			ขอ
29	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	มาเป็นครอบครัวใหญ่ ไร่หัวหมากน้อย ชื่อเป็นที่พัก (หาดนี้พักมาที่นี้ 3 ครั้งแล้ว) วิวสวยมาก	เพิ่มเติม			
30	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	เป็นรีสอร์ทใหญ่แบ่งออกเป็นหลายโซน ห้องนอนแบบวิวทะเลดีมาก วิวสวยมาก วิวสวยมาก วิวสวยมาก	เพิ่มเติม			
31	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	สิ่งสำคัญที่สุดของรีสอร์ทคือการบริการ ที่นี่ย่อย ดีมาก วิวในร่มมีต้นไม้ทั้ง น้อย พน. (ผู้)	เพิ่มเติม			ขอ
32	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png	"เกาะภูเก็ต" Koh Kood Resort เป็นที่พักที่สวยงาม และมีความเป็นส่วนตัว ...	เพิ่มเติม			ขอ
33	3/wwww.gstatic.com/images/branding/product/1x/googleleg_48dp.png					

ภาพประกอบ 5 ตัวอย่างข้อมูลหลังจาก scraping data

ในการคัดเลือกข้อมูลเพื่อให้การวิเคราะห์มีความถูกต้องและสอดคล้องกับวัตถุประสงค์ของการศึกษา ข้อมูลรีวิวที่ใช้ในการวิจัยจะต้องผ่านกระบวนการคัดเลือกตามเกณฑ์ที่กำหนดไว้ ดังนี้

1. **สถานที่ตั้งของแหล่งข้อมูล** ต้องเป็นข้อมูลรีวิวที่มาจากโรงแรมหรือที่พักซึ่งตั้งอยู่บนพื้นที่ของเกาะภูเก็ต จังหวัดตราด เท่านั้น เพื่อจำกัดขอบเขตของการวิเคราะห์ให้อยู่ภายในพื้นที่ศึกษาอย่างชัดเจน

2. **แหล่งที่มาของความคิดเห็น** ต้องเป็นรีวิวที่เขียนโดยผู้ใช้งานจริง (real user-generated content) ซึ่งมีการแสดงข้อความแสดงความคิดเห็นอย่างชัดเจน ไม่เป็นข้อความที่สร้างขึ้นโดยระบบหรือรีวิวจำลอง

3. **เนื้อหาของรีวิว** ต้องเป็นข้อที่เป็นเนื้อหาสำคัญ โดยตัดข้อความที่ไม่สื่อสารเนื้อหาเชิงความเห็นหรือมีเนื้อหาที่ไม่เกี่ยวข้อง เช่น ข้อความที่มีเพียงอีโมจิ การพิมพ์ซ้ำตัวอักษร หรือข้อความที่สั้นเกินไปและไม่ให้ข้อมูลที่เป็นประโยชน์ต่อการวิเคราะห์

4. **ภาษาที่ใช้ในการแสดงความคิดเห็น** ต้องเป็นข้อมูลรีวิวที่เขียนด้วยภาษาไทย เพื่อให้สอดคล้องกับการประมวลผลภาษาในแบบจำลองที่พัฒนาและเพื่อหลีกเลี่ยงปัญหาทางภาษาที่อาจส่งผลต่อความแม่นยำของการวิเคราะห์

ทั้งนี้ ข้อมูลรีวิวจากแพลตฟอร์ม Google Maps เป็นแหล่งข้อมูลหลักที่น่าสนใจสำหรับการนำมาใช้วิเคราะห์ครั้งนี้ เนื่องจากเป็นแหล่งที่มี ความหลากหลายและน่าเชื่อถือสูง โดยเปิดให้ผู้ใช้งานจริงสามารถแสดงความคิดเห็นเกี่ยวกับประสบการณ์ในการเข้าพักได้อย่างอิสระและต่อเนื่องในเชิงเวลา ข้อมูลที่ได้รับจึงมีลักษณะเป็น ความคิดเห็นที่สะท้อนอารมณ์ ความพึงพอใจ หรือข้อร้องเรียนอย่าง

ตรงไปตรงมา ซึ่งเป็นปัจจัยสำคัญที่เหมาะสมต่อการนำมาวิเคราะห์ด้วยเทคนิค การวิเคราะห์ความรู้สึก (Sentiment Analysis) โดยเฉพาะเมื่อมีการคัดเลือกข้อมูลตามเกณฑ์ที่เข้มงวด เช่น การเลือกเฉพาะรีวิวที่มาจากที่พักบนเกาะกูด เป็นภาษาไทย และมีเนื้อหาสาระในเชิงความเห็นอย่างชัดเจน

ข้อมูลที่ผ่านกระบวนการคัดเลือกแล้วจะถูกนำมาผ่านการวิเคราะห์และตรวจสอบโครงสร้างของข้อมูล (Structural Analysis) ดังภาพประกอบ 6 เพื่อทำความเข้าใจลักษณะของข้อมูลที่ใช้ในการวิจัยอย่างเป็นระบบ ดังนี้

1. การแยกหน่วยข้อมูล (Data Record Identification) แยกรีวิวแต่ละรายการให้อยู่ในรูปแบบของข้อมูลเชิงรายการ (record) โดยประกอบด้วยองค์ประกอบหลัก เช่น ข้อความรีวิว คะแนนความพึงพอใจ วันที่โพสต์ ชื่อที่พัก และตำแหน่งที่ตั้ง

2. การจำแนกประเภทของข้อมูล (Variable Classification) จัดประเภทข้อมูลในแต่ละฟิลด์ เช่น

- ข้อความรีวิว → ข้อมูลไม่เป็นโครงสร้าง (Unstructured Text)
- คะแนนรีวิว → ข้อมูลเชิงตัวเลข (Numerical Data)
- วันที่โพสต์ → ข้อมูลเชิงเวลา (Temporal Data)
- ชื่อสถานที่ → ข้อมูลเชิงหมวดหมู่ (Categorical Data)

3. การตรวจสอบความสมบูรณ์ของข้อมูล (Data Completeness Check) ตรวจสอบว่าข้อมูลแต่ละรีวิวมีองค์ประกอบครบถ้วน เช่น มีข้อความรีวิว มีคะแนน และมีวันที่โพสต์ โดยตัดรายการที่ขาดข้อมูลสำคัญหรือไม่ตรงตามเกณฑ์ที่กำหนด

4. การคัดกรองข้อมูลที่ไม่มีความเกี่ยวข้อง (Content Relevance Filtering) กำจัดรีวิวที่ไม่มีเนื้อหา เช่น ข้อความที่มีเพียงอีโมจิ คำซ้ำสั้น ๆ หรือข้อความที่ไม่สะท้อนความคิดเห็นที่แท้จริง

5. การตรวจสอบภาษา (Language Detection and Filtering) คัดเลือกเฉพาะรีวิวที่เขียนด้วยภาษาไทย โดยอาจยกเว้นในกรณีที่ใช้คำทับศัพท์ภาษาอังกฤษ

6. การจัดรูปแบบข้อมูล (Data Structuring) นำข้อมูลที่ได้จากการคัดกรองมาจัดให้อยู่ในรูปแบบตาราง (structured tabular format) เพื่อความสะดวกในการนำไปใช้งานในขั้นตอนถัดไป เช่น การประมวลผลข้อความ การแปลงข้อมูลเป็นเวกเตอร์ และการสร้างแบบจำลองการวิเคราะห์ความรู้สึก

A	C	E	H	I		
Rate	Standard	Time Ago	no	File Paths	Noise Removed	Sentiment
2	11 months ago on	Review_KohKoodParadiseBeachResort.csv		Reception ควระเอกที่ฟากกว่าห้าจะไม่ใช้ว่าตามแล้วแต่ไม่มีแอลกอฮอล์อะไร การรับแขกตอนเช็คอินก็ดี ไม่ค่อย	negative	
5	11 months ago on	Review_TinkerbellResort.csv		Right บนชายหาด เข้าตัวและพระอาทิตย์ตกดิน time ที่ดีที่สุดสำหรับการว่ายน้ำ เมื่อรวมกับ Netflix ไว้ด้วย	positive	
5	10 months ago on	Review_HighSeasonPoolVilla.csv		Service ดี พนักงานบริการดี น่ารัก โรงแรมก็ดีด้วย โดยรวมแล้วประทับใจมาก	positive	
5	11 months ago on	Review_KoKutAoPhraoBeachResort.csv		เกาะภูเก็ต หรือ เกาะภูเก็ต ห่างจากกรุงเทพ 350 กิโลเมตร เกาะมีถนนเส้นเดียวโดยเจตนาเพื่อไม่ให้ทำด้วยธรรม	positive	
5	11 months ago on	Review_HideoutKohKood.csv		เข้าพักSeaview pool villa 1 คืน ห้องพักสวยงาม เหมือนหลุดออกมาจาก pinterest พนักงานพร้อมให้บริการดี	positive	
5	11 months ago on	Review_S-beachResort.csv		เจ้าของเป็นกันเอง พนักงานบริการดีมาก ห้องพักสะอาด อาหารอร่อยทุกอย่าง โดยเฉพาะหมักผัดไข่เค็ม หวาน	positive	
5	11 months ago on	Review_WoodenHutKohKood.csv		เจ้านายเป็นมิตรมากและทำให้เรารู้สึกเหมือนอยู่บ้าน เขาจัดงานเลี้ยงอาหารค่ำกับแขกทุกคนที่เขาทำอาหาร	positive	
5	10 months ago on	Review_WoodenHutKohKood.csv		เจ้าภาพทำอาหารเพื่ออาสาทุกคนดูแลความต้องการของทุกคนและช่วยเหลือดีมาก นักรักเหมือนอยู่บ้าน	positive	
5	11 months ago on	Review_CaptainHookResort.csv		เป็นที่พักเหมาะกับการพักผ่อน พนักงานบริการดีมาก มากครั้งประทับใจ	positive	
5	11 months ago on	Review_CaptainHookResort.csv		เป็นประสบการณ์ที่ดี กิจกรรมเยอะ สนุก พนักงานบริการดี บรรยากาศดี อาหารดี ค่ำดีมาก	positive	
5	11 months ago on	Review_PeterPanResort.csv		เป็นสถานที่ที่ยังคงมีความเป็นธรรมชาติ ไม่วุ่นวาย เหมาะกับการพักผ่อนแบบ slowlife เราจองมาแบบแพคเกจ	positive	
3	11 months ago on	Review_TheMermaidHouse.csv		เมื่อมาถึงแรกต้องให้เราได้รับการต้อนรับด้วยใบหน้าที่ไม่เป็นมิตรของเจ้าของผู้หญิง ฉะนั้นอีก 4 คนใน	neutral	
5	10 months ago on	Review_HighSeasonPoolVilla.csv		เราใช้เวลา 3 คืนที่ไอซ์แลนด์และการเข้าพักกันทั้งหมวก เกาะนี้เงียบสงบมากและมันก็ท่องเที่ยวด้วย ชายหาดมี	positive	
5	11 months ago on	Review_HighSeasonPoolVilla.csv		เราค้นพบรีสอร์ทแห่งใหม่เมื่อไม่กี่ปีก่อนในฐานะคู่รักและกลับมาเป็นครอบครัว เป็นเพียงประสบการณ์ที่ยอดเยี่ยม	positive	
2	10 months ago on	Review_SecretKohKood.csv		เราเข้าพักที่นี่กับครอบครัว ที่เลือกที่นี่เพราะดูรีวิวจากหลายๆ ที่ว่าสวย พอมาถึง ที่นี่ก็สวยจริงค่ะ บ้านพักดี	negative	
5	11 months ago on	Review_EveHouse.csv		เรามีการเข้าพักที่ยอดเยี่ยมนึงสัปดาห์ในไม่ช้าก็ไปนอนที่บ้านของอีฟ บังคับให้สะดวกสบายตั้งอยู่ในสวนที่สวยงาม	positive	
5	10 months ago on	Review_HighSeasonPoolVilla.csv		เรามีวิลล่าพร้อมสระว่ายน้ำ พนักงานทุกคนสุภาพ และทำให้เรารู้สึกเหมือนอยู่ที่นี่ เราทำมาตั้งแต่จาวา	positive	
3	11 months ago on	Review_KohKoodParadiseBeachResort.csv		เราขึ้นอยู่กับเวลา 3 วัน มันยอดเยี่ยมจริงๆ แต่ห้องต้องได้รับการปรับปรุงหลังปลูกเล็กน้อย เมื่อเสร็จ	neutral	
3	11 months ago on	Review_KohKoodBeachResort.csv		เราอยู่ที่นั่นสองคืน แต่เมื่อเรากำลังจะออกไป พวกเราตรวจสอบห้องและบอกว่ามีถุงซักผ้าหาย และพวกเร	neutral	
2	11 months ago on	Review_AoNoiResort.csv		เห็นด้วยกับรีวิวด้านล่าง ทุกประการ	negative	
1	11 months ago on	Review_SuanyKohKood.csv		แอลกอฮอล์ในโรงเบียร์อยู่อย่าง เยอะ	negative	
5	11 months ago on	Review_SuanyKohKood.csv		โดยรวมชอบมากมีที่พักที่มีปลาทะเลให้ดูแล้วค่ะ	positive	
5	11 months ago on	Review_KoKutAoPhraoBeachResort.csv		โดยรวมดีเกินคาด อาหารอร่อย ทะเลสวย บริการดีเยี่ยมค่ะ นอนหลับสบาย	positive	
2	11 months ago on	Review_PeterPanResort.csv		โรงแรม 3 ดาว สิ่งอำนวยความสะดวกอย่าง2ดาว ราคาส่งทางเข้ารีสอร์ทที่ไม่มีต้องเดินไปเอาเอง	negative	
5	10 months ago on	Review_S-beachResort.csv		โรงแรมเกินความคาดหมาย เราได้รับการที่ทักทายอย่างน่าอัศจรรย์ พื้นที่เขียวสวย เป็นระเบียบเรียบร้อย บ้านอยู่	positive	
1	10 months ago on	Review_HideoutKohKood.csv		โรงแรมเป็นการพักผ่อนที่ดี ทรัพยากรและยอดเยี่ยมเหนือ Google ดึง Facebook และหมายเลขโทรศัพท์ เราเอง	negative	
5	11 months ago on	Review_TinkerbellResort.csv		โรงแรมที่ดีมากจนหาที่พักที่สวยที่สุดแห่งหนึ่งบนเกาะ	positive	

ภาพประกอบ 6 ตัวอย่างข้อมูลที่ได้จากการวิเคราะห์โครงสร้างข้อมูล (Structural Analysis)

ข้อมูลที่ได้ผ่านกระบวนการวิเคราะห์และตรวจสอบโครงสร้างของข้อมูลจะถูกนำมาวิเคราะห์ข้อความ (Textual Analysis) (ภาพประกอบ 7) ซึ่งใช้ในการสกัดคุณลักษณะเชิงภาษาและความหมายจากข้อมูลรีวิวในรูปแบบข้อความไม่เป็นโครงสร้าง (unstructured text) โดยมีขั้นตอนในการดำเนินการดังนี้

1. การทำความสะอาดข้อความ (Text Cleaning) ข้อความรีวิวจะถูกทำความสะอาดเพื่อลดสิ่งรบกวนที่ไม่จำเป็น เช่น การลบอักขระพิเศษ สัญลักษณ์ อีโมจิ URL ตัวเลข และการแปลงข้อความให้เป็นตัวพิมพ์เดียวกัน (lowercase) ซึ่งช่วยให้การประมวลผลภาษาเป็นไปอย่างมีประสิทธิภาพ

2. การตัดคำภาษาไทย (Word Tokenization) ใช้เครื่องมือตัดคำที่เหมาะสมกับภาษาไทย เช่น PyThaiNLP โดยเฉพาะ engine อย่าง attacut หรือ newmm เพื่อแบ่งข้อความออกเป็นคำหรือหน่วยย่อยที่มีความหมาย ซึ่งเป็นพื้นฐานสำคัญในการสร้างเวกเตอร์หรือลำดับคำสำหรับโมเดล

3. การตัดคำฟุ่มเฟือย (Stopword Removal) ลบคำที่ไม่สื่อความหมายเชิงวิเคราะห์ เช่น คำว่า “คือ” “ของ” “และ” “ก็” “ว่า” โดยอ้างอิงจากรายการคำหยุด (stopword list) ของภาษาไทย เพื่อให้เหลือเฉพาะคำที่มีนัยสำคัญในการวิเคราะห์ความรู้สึก

4. การแปลงข้อความเป็นเวกเตอร์ (Text Vectorization) นำข้อความที่ได้จากการตัดคำมาแปลงให้อยู่ในรูปแบบเวกเตอร์เพื่อให้สามารถป้อนเข้าสู่แบบจำลองการเรียนรู้ได้ อาทิ Word Embedding เช่น FastText สำหรับโมเดล Deep Learning และ Tokenizer จากโมเดล Pretrained เช่น BERT Tokenizer สำหรับการเรียนรู้เชิงบริบท

5. การแปลง Label (Label Encoding) ปรับป้ายกำกับของข้อความ (Sentiment label) ให้อยู่ในรูปแบบตัวเลข หรือแบบ One-hot encoding เช่น positive = [1,0,0], neutral = [0,1,0], negative = [0,0,1] เพื่อให้สามารถใช้กับโมเดลการจำแนกประเภทได้

6. การตรวจสอบความสมบูรณ์ของข้อมูล (Integrity Check) ตรวจสอบความครบถ้วนของข้อความหลังจากการประมวลผล เช่น ข้อความที่เหลือคำไว้น้อยเกินไปจะถูกกรองออก เพื่อให้มั่นใจว่าข้อมูลมีความพร้อมในการฝึกโมเดล



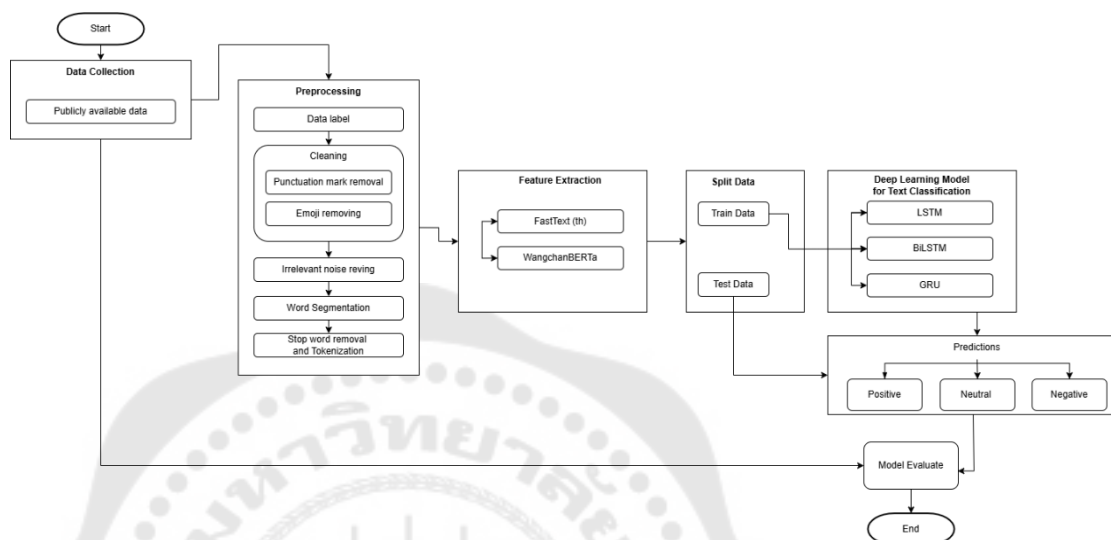
ภาพประกอบ 7 ตัวอย่างข้อมูลที่ผ่านการวิเคราะห์ข้อความ (Textual Analysis)

โดยแสดงผลในรูปแบบ Word cloud

3.3 กระบวนการทำงานของแบบจำลอง

กระบวนการทำงานของแบบจำลองครอบคลุมขั้นตอนสำคัญในการวิเคราะห์ความรู้สึกจากข้อความภาษาไทย ตั้งแต่การเก็บรวบรวมข้อมูลจากแหล่งสาธารณะ การทำความสะอาดและเตรียม

ข้อมูล การแปลงข้อความเป็นเวกเตอร์คุณลักษณะ (Feature Vector) จนถึงการสร้างแบบจำลองเชิงลำดับ และการประเมินผลลัพธ์ที่ได้จากการจำแนกประเภทข้อความ ดังแสดงในภาพประกอบ 8



ภาพประกอบ 8 ขั้นตอนการทำงานของแบบจำลอง

ขั้นตอนการดำเนินงานของแบบจำลองเริ่มจากกระบวนการเก็บรวบรวมข้อมูล โดยคัดเลือกความคิดเห็นของผู้ใช้งานที่แสดงต่อที่พักและการให้บริการของโรงแรมบนเกาะกูด ผ่านแพลตฟอร์ม Google Maps ซึ่งเป็นแหล่งข้อมูลสาธารณะที่เปิดให้ผู้ใช้งานสามารถแสดงความคิดเห็นได้โดยตรง ชุดข้อมูลที่ได้ประกอบด้วยวิธีวิที่สะท้อนประสบการณ์จริงของผู้เข้าพัก และมีความหลากหลายทั้งในเชิงเนื้อหาและอารมณ์ ทำให้เหมาะสมอย่างยิ่งในการนำมาใช้เป็นข้อมูลขาเข้า (input data) สำหรับการวิเคราะห์ความรู้สึก (Sentiment Analysis)

ภายหลังการรวบรวมข้อมูล (Data Collection) ได้ดำเนินการทำความสะอาดข้อมูล (Preprocessing) เพื่อปรับปรุงคุณภาพของข้อมูลก่อนนำเข้าสู่กระบวนการประมวลผล โดยการลบองค์ประกอบที่ไม่จำเป็น เช่น HTML Tag, URL, สัญลักษณ์เครื่องหมายคำพูด และอักขระพิเศษอื่น ๆ ที่อาจรบกวนการวิเคราะห์ จากนั้นดำเนินการตรวจสอบและแก้ไขคำสะกดผิดในข้อความรีวิว ด้วยเครื่องมือตรวจสอบการสะกดคำ (Spell Checker) จากไลบรารี PyThaiNLP ซึ่งช่วยให้ข้อมูลมีความถูกต้องมากยิ่งขึ้น และมีความพร้อมสำหรับการประมวลผลในขั้นตอนการสร้างแบบจำลอง

หลังจากขั้นตอนการทำความสะอาดข้อมูล ข้อมูลข้อความจะถูกแปลงให้อยู่ในรูปแบบของเวกเตอร์คุณลักษณะ (Feature Vector) เพื่อให้สามารถนำไปใช้เป็นข้อมูลนำเข้าในการฝึกแบบจำลอง โดยประยุกต์ใช้เทคนิคการแปลงข้อความเป็นเวกเตอร์จำนวน 2 รูปแบบ ได้แก่

1. FastText (th) ซึ่งสามารถวิเคราะห์คำจาก subword ได้ดีและรองรับภาษาไทย
2. WangchanBERTa ซึ่งเป็นโมเดลภาษาที่ฝึกจากข้อมูลภาษาไทยโดยเฉพาะ และสามารถเรียนรู้บริบทในระดับลึกได้ดี

จากนั้นได้ดำเนินการแบ่งชุดข้อมูลออกเป็นสองส่วน ได้แก่ ชุดข้อมูลสำหรับการฝึกสอน (training set) ร้อยละ 80 และชุดข้อมูลสำหรับการทดสอบ (testing set) ร้อยละ 20 เพื่อใช้ในการประเมินประสิทธิภาพของแบบจำลอง โดยในการศึกษาได้นำชุดข้อมูลเดียวกันไปใช้ในการฝึกแบบจำลองจำนวน 3 เทคนิค ได้แก่ LSTM (Long Short-Term Memory), BiLSTM (Bidirectional LSTM) และ GRU (Gated Recurrent Unit) ซึ่งเป็นแบบจำลองโครงข่ายประสาทเทียมเชิงลำดับ (Sequential Neural Networks) ที่ได้รับความนิยมในงานประมวลผลภาษาธรรมชาติ (Natural Language Processing) หลังจากแบบจำลองแต่ละเทคนิคได้รับการฝึกฝนเรียบร้อยแล้ว จะนำไปใช้กับชุดข้อมูลทดสอบเพื่อทำการพยากรณ์ประเภทของความรู้สึกในข้อความรีวิว (เชิงบวก เป็นกลาง หรือเชิงลบ) และประเมินประสิทธิภาพของแบบจำลองด้วยค่าชี้วัดมาตรฐาน ได้แก่ Accuracy, Precision, Recall และ F1-score เพื่อใช้เปรียบเทียบผลลัพธ์ระหว่างแบบจำลอง และพิจารณาเลือกเทคนิคที่ให้ผลลัพธ์ที่ดีที่สุดในบริบทของการวิเคราะห์ความรู้สึกจากรีวิวโรงแรมที่เขียนเป็นภาษาไทย

3.4 การเตรียมข้อมูล

ในการเตรียมข้อมูลสำหรับการวิเคราะห์ความรู้สึกจากข้อมูลรีวิวโรงแรมซึ่งจัดเก็บในรูปแบบไฟล์ CSV โดยกำหนดให้แถวแรกของไฟล์เป็นชื่อของแต่ละข้อมูล จากนั้นแปลงข้อมูลให้อยู่ในรูปแบบ DataFrame เพื่อความสะดวกในการจัดการข้อมูลในขั้นตอนต่าง ๆ และปรับปรุงชื่อหัวข้อของแต่ละคอลัมน์ให้สอดคล้องกับลักษณะของข้อมูลที่แท้จริง เพื่อความชัดเจนและเหมาะสมสำหรับขั้นตอนการประมวลผลข้อมูล โดยขั้นตอนการเตรียมข้อมูลประกอบด้วย

1. การกำหนดป้ายกำกับข้อมูล (Data Labeling) ทำการจัดประเภทความคิดเห็นตามอารมณ์ของผู้รีวิวออกเป็น 3 กลุ่ม ได้แก่

- ความรู้สึกเชิงบวก โดยใช้ข้อความรีวิวที่ผู้ให้บริการให้คะแนน Rate = 5 มีจำนวน 3,896 ข้อความ
- ความรู้สึกเชิงลบ โดยใช้ข้อความรีวิวที่ผู้ให้บริการให้คะแนน Rate = 1 มีจำนวน 1,000 ข้อความ
- ความรู้สึกเป็นกลาง โดยใช้ข้อความรีวิวที่ผู้ให้บริการให้คะแนน Rate = 3 มีจำนวน 1,279 ข้อความ

จากนั้นทำ under sampling โดยอ้างอิงจาก Rate = 1 ที่เก็บข้อมูลมาได้น้อยที่สุดจำนวน 1,000 ข้อความ เพื่อให้ข้อมูล balance กัน เพื่อใช้เป็นตัวแปรเป้าหมาย (Target Variable) สำหรับการเรียนรู้ของแบบจำลอง และในส่วน Rate = 2 และ 4 นั้นไม่ได้เลือกเข้ามาเป็นจำนวน โดย Rate = 2 มีจำนวน 253 ข้อความ และ Rate = 4 มีจำนวน 1,683 ข้อความ

2. ทำความสะอาดข้อมูล (Data Cleaning) ลบองค์ประกอบที่ไม่จำเป็นซึ่งอาจรบกวนการประมวลผล เช่น HTML Tag, URL, อีโมจิ และเครื่องหมายวรรคตอนที่ไม่มีความสำคัญต่อความหมายของข้อความ

3. การลบข้อความรบกวน (Noise Removal) ดำเนินการกรองข้อความที่ซ้ำซ้อน หรือมีลักษณะไม่เกี่ยวข้องกับเนื้อหาหลัก เช่น คำซ้ำ คำหยาบ หรือถ้อยคำที่ไม่สะท้อนความคิดเห็นจริงของผู้ใช้

4. การตัดคำ (Word Segmentation) สำหรับข้อความภาษาไทย ซึ่งไม่มีการเว้นวรรคระหว่างคำ จำเป็นต้องใช้กระบวนการตัดคำเพื่อแบ่งข้อความออกเป็นหน่วยคำ (token) ที่สามารถนำไปใช้วิเคราะห์ได้ อย่างไรก็ตาม สำหรับข้อความที่ใช้กับแบบจำลอง WangchanBERTa ไม่จำเป็นต้องดำเนินการตัดคำในขั้นตอนนี้ เนื่องจากตัวแบบมีระบบการตัดคำในตัว (built-in tokenizer) ที่รองรับการประมวลผลภาษาไทยโดยเฉพาะอยู่แล้ว

5. การลบคำฟุ่มเฟือย (Stopword Removal) และการแปลงหน่วยคำ (Tokenization) ดำเนินการลบคำที่ไม่ก่อให้เกิดนัยสำคัญทางการวิเคราะห์ เช่น คำบุพบท คำสันธาน และคำซ้ำทั่วไป พร้อมทั้งแปลงข้อความให้เป็นลำดับของหน่วยคำ (tokens) เพื่อเตรียมสำหรับการนำเข้าสู่กระบวนการสร้างแบบจำลองในขั้นตอนถัดไป

ในการวิจัยครั้งนี้ ได้ดำเนินการวิเคราะห์ข้อมูลเบื้องต้น เพื่อทำความเข้าใจลักษณะและโครงสร้างของชุดข้อมูลรีวิวก่อนใช้ในการจำแนกความรู้สึก โดยเริ่มจากการตรวจสอบการกระจายตัวของแต่ละคลาส (เชิงบวก เป็นกลาง เชิงลบ) เพื่อประเมินความสมดุลของข้อมูล จากนั้นวิเคราะห์ความยาวของข้อความในแต่ละคลาสเพื่อกำหนดความยาวสูงสุดที่เหมาะสมสำหรับการฝังข้อความ (text embedding) พร้อมทั้งตรวจสอบคำที่พบมากที่สุดในการรวมและในแต่ละคลาส เพื่อใช้ในการคัดเลือกคำสำคัญและวางแผนการทำ Feature Engineering นอกจากนี้ยังแสดงภาพรวมคำที่พบบ่อยผ่าน Word Cloud และวิเคราะห์คำเชิงลำดับ (n-gram) เพื่อศึกษารูปแบบการใช้ภาษาที่เกี่ยวข้องกับอารมณ์หรือความรู้สึกในข้อความ ทั้งนี้กระบวนการ EDA มีบทบาทสำคัญในการเตรียมข้อมูลให้เหมาะสมกับการพัฒนาแบบจำลองวิเคราะห์ความรู้สึกอย่างมีประสิทธิภาพในขั้นตอนต่อไป

Negative	คำ 1 คำ	จำนวนครั้ง	คำ 2 คำ	จำนวนครั้ง	คำ 3 คำ	จำนวนครั้ง
1	ห้อง	318	(โรง, แรม)	114	(ขาย, หาด, สบายงาม)	12
2	พัก	249	(ขาย, หาด)	104	(ผ้า, เช็ด, ตัว)	11
3	ดี	227	(ห้อง, พัก)	82	(ห้อง, พัก, เก่า)	8
4	อาหาร	174	(ร้าน, อาหาร)	50	(จอง, ห้อง, พัก)	7
5	แย	133	(อาหาร, เข้า)	37	(พัก, โรง, แรม)	7
6	ขาย	123	(ห้อง, น้ำ)	33	(อาหาร, เข้า, แยม)	7
7	พนักงาน	122	(แหม่ม, ต้อนรับ)	26	(โรง, แรม, โรง)	6
8	โรง	118	(บริการ, แยม)	20	(แรม, โรง, แรม)	6
9	หาด	115	(ราคา, แพง)	19	(โรง, แรม, แยม)	6
10	แรม	114	(บท, วิจารณ์)	18	(ขาย, หาด, ตัว)	6
11	บริการ	102	(จอง, ห้อง)	17	(อาหาร, เข้า, ดี)	6
12	จอง	94	(กลิ่น, เหม็น)	14	(ผ้า, ปู, นอน)	6
13	กิน	94	(ดู, เหมือน)	13	(ห้อง, พัก, ดี)	5
14	คน	78	(พัก, กิน)	13	(ร้าน, อาหาร, ดี)	5
15	ราคา	76	(ดูแล, ดี)	12	(ดี, ห้อง, พัก)	5
16	รีสอร์ท	76	(พัก, ดี)	12	(ริม, ขาย, หาด)	5
17	เหมือน	65	(บรรยากาศ, ดี)	12	(ขาย, หาด, สาธารณะ)	5
18	ตัว	62	(ดี, ห้อง)	12	(บท, วิจารณ์, เชิง)	5
19	สถานที่	58	(พนักงาน, มืด)	12	(วิจารณ์, เชิง, บวก)	5
20	ต้อนรับ	57	(ภาษา, อังกฤษ)	12	(จอง, โรง, แรม)	5

ภาพประกอบ 9 n-gram ในความคิดเห็นเชิงลบ

จากภาพประกอบ 9 คือ ตัวอย่าง n-gram ของความคิดเห็นเชิงลบจำนวน 20 ลำดับแรกที่พบบ่อยที่สุด

Neutral	คำ 1 คำ	จำนวนครั้ง	คำ 2 คำ	จำนวนครั้ง	คำ 3 คำ	จำนวนครั้ง
1	ดี	486	(ขาย, หาด)	314	(ขาย, หาด, สบายงาม)	60
2	ขายหาด	343	(ห้อง, พัก)	191	(ขาย, หาด, ดี)	26
3	ห้อง	307	(โรง, แรม)	176	(ขาย, หาด, สบาย)	22
4	น้ำ	213	(อาหาร, เข้า)	118	(ริม, ขาย, หาด)	20
5	รีสอร์ท	201	(ร้าน, อาหาร)	100	(โรง, แรม, ดี)	19
6	ห้องพัก	194	(หาด, สบายงาม)	60	(ดี, ขาย, หาด)	18
7	โรงแรม	189	(ห้อง, น้ำ)	56	(ห้อง, พัก, สะอาด)	17
8	พนักงาน	183	(เจียน, สงบ)	47	(ห้อง, พัก, ดี)	17
9	พัก	180	(บริการ, ดี)	41	(อาหาร, เข้า, ดี)	14
10	อาหาร	170	(สระ, ว่ายน้ำ)	39	(ดี, ห้อง, พัก)	13
11	สวยงาม	169	(หาด, สบาย)	38	(พนักงาน, ช่วยเหลือ, ดี)	12
12	หรับ	153	(ราคา, แพง)	34	(ดี, อาหาร, เข้า)	11
13	สวย	145	(ดี, อาหาร)	34	(หาด, สวย, เกะ)	11
14	สำ	143	(อาหาร, อร่อย)	33	(โรง, แรม, ขาย)	11
15	บังกะไล	129	(ทำเล, ดี)	32	(พนักงาน, บริการ, ดี)	11
16	เกาะ	126	(ดี, พนักงาน)	30	(ขาย, หาด, ตัว)	11
17	ราคา	119	(พัก, ดี)	29	(ห้อง, พัก, โอเค)	10
18	อาหารเข้า	118	(ดี, ห้อง)	28	(แรม, ขาย, หาด)	10
19	สถานที่	108	(หาด, ดี)	27	(สระ, ว่ายน้ำ, น้ำ)	10
20	บริการ	107	(บรรยากาศ, ดี)	27	(อาหาร, ร้าน, อาหาร)	10

ภาพประกอบ 10 n-gram ในความคิดเห็นเป็นกลาง

จากภาพประกอบ 10 คือ ตัวอย่าง n-gram ของความคิดเห็นเป็นกลางจำนวน 20 ลำดับแรกที่พบบ่อยที่สุด

Positive	คำ 1 คำ	จำนวนครั้ง	คำ 2 คำ	จำนวนครั้ง	คำ 3 คำ	จำนวนครั้ง
1	ดี	751	(ขาย, หาด)	237	(ห้อง, พัก, สะอาด)	55
2	พัก	616	(ห้อง, พัก)	208	(พนักงาน, บริการ, ดี)	39
3	อาหาร	448	(โรง, แรม)	159	(ขาย, หาด, สวยงาม)	30
4	ห้อง	358	(บริการ, ดี)	109	(ริม, ขาย, หาด)	27
5	หาด	339	(อาหาร, เข้า)	105	(ขาย, หาด, สวย)	25
6	พนักงาน	317	(อาหาร, อร่อย)	101	(อาหาร, เข้า, อร่อย)	22
7	ขาย	265	(เจียบ, สงบ)	97	(โรง, แรม, ดี)	21
8	บริการ	239	(ร้าน, อาหาร)	88	(ดี, ห้อง, พัก)	20
9	สะอาด	205	(พัก, สะอาด)	64	(ห้อง, พัก, ดี)	20
10	สวย	197	(ช่วยเหลือ, ดี)	58	(ดี, อาหาร, อร่อย)	18
11	ยอดเยี่ยม	190	(หาด, สวย)	56	(ดี, ขาย, หาด)	18
12	สวยงาม	183	(บรรยากาศ, ดี)	55	(ดี, อาหาร, เข้า)	16
13	โรง	170	(พนักงาน, บริการ)	50	(เรือ, คา, บัค)	16
14	อร่อย	165	(ดี, อาหาร)	50	(พนักงาน, ช่วยเหลือ, ดี)	16
15	แรม	159	(พัก, ดี)	48	(บริการ, ดี, อาหาร)	15
16	ถนน	157	(ดูแล, ดี)	38	(สระ, ว่าย, น้ำ)	15
17	สถานที่	151	(สระ, ว่ายน้ำ)	34	(อาหาร, เข้า, ดี)	15
18	รีสอร์ท	137	(พระอาทิตย์, ตก)	32	(ขาย, หาด, ตัว)	14
19	เกาะ	137	(เกาะ, กุด)	32	(หาด, สวย, เกาะ)	13
20	รัก	127	(หาด, สวยงาม)	32	(ห้อง, พัก, สวย)	12

ภาพประกอบ 11 n-gram ในความคิดเห็นเชิงบวก

จากภาพประกอบ 11 คือ ตัวอย่าง n-gram ของความคิดเห็นเชิงบวกจำนวน 20 ลำดับแรกที่พบบ่อยที่สุด

3.5 การสร้างแบบจำลอง

ในงานวิจัยนี้ ก่อนจะนำข้อความเข้าสู่กระบวนการฝึกโมเดล จะต้องแปลงข้อมูลให้อยู่ในรูปแบบตัวเลข (Vectorization) โดยใช้เทคนิค Word Embedding ที่สามารถรักษาความหมายและบริบทของคำในภาษาไทยได้อย่างเหมาะสม ซึ่งในการวิจัยนี้ได้เลือกใช้ 2 เทคนิคหลัก ได้แก่

1. FastText (th): เทคนิคการฝังคำที่ใช้ข้อมูล subword ทำให้สามารถวิเคราะห์คำที่ไม่อยู่ในพจนานุกรม (Out-of-Vocabulary) และคำเฉพาะในภาษาไทยได้อย่างมีประสิทธิภาพ

2. WangchanBERTa: โมเดลประมวลผลภาษาที่ถูกฝึกด้วยข้อมูลภาษาไทยขนาดใหญ่ สามารถเรียนรู้ความสัมพันธ์และบริบทของคำในประโยคได้ดี โดยใช้สถาปัตยกรรมของ RoBERTa ซึ่งเหมาะกับการทำงานเชิงลึก (Deep Contextual Understanding)

หลังจากได้เวกเตอร์คุณลักษณะที่เหมาะสมแล้ว ข้อมูลจะถูกส่งเข้าสู่กระบวนการฝึกเพื่อทดสอบแบบจำลองที่เหมาะสมกับชุดข้อมูล โดยใช้เทคนิค FastText (th) หรือ WangchanBERTa ร่วมกับแบบจำลอง LSTM, BiLSTM และ GRU ซึ่งสามารถออกแบบแบบจำลองได้ทั้งหมด 5 รูปแบบดังนี้

1. WangchanBERTa + LSTM เป็นการนำเทคนิคการฝังข้อความ (Text Embedding) ที่ใช้โมเดล WangchanBERTa ใช้ร่วมกับ LSTM (Long Short-Term Memory) ที่มีความสามารถในการจดจำบริบทของข้อมูลที่มีลำดับ เช่น ข้อความหรือประโยค โดย LSTM สามารถจดจำลำดับก่อนหน้า

และบริบทระยะยาวได้ดีกว่าโครงข่ายประสาทเทียมทั่วไป จึงเหมาะสำหรับการจำแนกอารมณ์หรือวิเคราะห์ความรู้สึกจากข้อความที่มีความซับซ้อน การใช้เทคนิคทั้งสองร่วมกันจึงเป็นการผสมผสานจุดแข็งของทั้งสองแนวทาง คือ ความสามารถในการเข้าใจความหมายเชิงบริบทในระดับคำจาก WangchanBERTa และความสามารถในการประมวลผลลำดับข้อมูลจาก LSTM ทำให้โมเดลนี้มีศักยภาพสูงในการวิเคราะห์ความรู้สึกจากข้อความภาษาไทย โดยเฉพาะในกรณีที่มีความหลากหลายของโครงสร้างภาษาและการใช้คำในบริบทจากผู้ใช้งานจริง

2. WangchanBERTa + BiLSTM เป็นการนำเทคนิคการฝังข้อความ (Text Embedding) ที่ใช้โมเดล WangchanBERTa ร่วมกับโครงข่ายประสาทเทียมแบบ BiLSTM (Bidirectional Long Short-Term Memory) ซึ่งมีความสามารถในการประมวลผลลำดับข้อมูลในสองทิศทาง คือทั้งจากอดีตไปอนาคต (forward) และจากอนาคตกลับสู่อดีต (backward) โดย BiLSTM ช่วยให้โมเดลเข้าใจบริบทของคำที่อยู่ทั้งก่อนหน้าและถัดไปในประโยค ซึ่งมีความสำคัญต่อการตีความอารมณ์หรือความหมายแฝงในข้อความที่มีความซับซ้อน การผสมผสานระหว่าง WangchanBERTa ซึ่งเชี่ยวชาญในการเข้าใจความหมายเชิงบริบทในระดับคำสำหรับภาษาไทย กับ BiLSTM ที่สามารถเรียนรู้ความสัมพันธ์เชิงลำดับจากทั้งสองทิศทาง จึงทำให้โมเดลนี้มีศักยภาพในการวิเคราะห์ความรู้สึกจากข้อความที่ใช้ภาษาหลากหลาย โครงสร้างซับซ้อน และมีลำดับคำที่ส่งผลต่อความหมาย เช่น ข้อความรีวิวโรงแรมที่มักใช้ภาษาพูดหรือการเปรียบเปรยที่ละเอียดอ่อน

3. FastText + LSTM เป็นการผสมผสานระหว่างเทคนิคการฝังข้อความด้วย FastText ซึ่งสามารถจับโครงสร้างของคำย่อย (subword) ได้ดี โดยเฉพาะในภาษาไทยที่มีลักษณะการรวมคำซับซ้อน และโครงข่ายประสาทเทียมแบบ LSTM (Long Short-Term Memory) ซึ่งมีความสามารถในการจดจำลำดับข้อมูลและบริบทระยะยาวได้ดีกว่าโครงข่ายทั่วไป การใช้ FastText ช่วยให้โมเดลเข้าใจความหมายของคำที่อาจไม่เคยพบมาก่อน (out-of-vocabulary) ได้บางส่วน ขณะที่ LSTM ช่วยวิเคราะห์ความสัมพันธ์ระหว่างคำตามลำดับในประโยค จึงเหมาะสำหรับการวิเคราะห์ความรู้สึกจากข้อความภาษาไทยที่ใช้ภาษาธรรมชาติและมีโครงสร้างยืดหยุ่น

4. FastText + BiLSTM เป็นการนำการฝังข้อความด้วย FastText ซึ่งมีจุดเด่นในการแยกคำย่อยและรักษาความหมายของคำใกล้เคียง มาใช้ร่วมกับโครงข่ายประสาทเทียมแบบ BiLSTM (Bidirectional Long Short-Term Memory) ที่สามารถประมวลผลลำดับคำได้ทั้งในทิศทางจากซ้ายไปขวาและจากขวาไปซ้าย ทำให้สามารถเข้าใจบริบทของคำได้ครบถ้วนมากยิ่งขึ้น โดยเฉพาะในกรณีที่ความหมายของคำใดคำหนึ่งขึ้นอยู่กับคำรอบข้างทั้งก่อนหน้าและถัดไป การรวมโมเดลทั้งสองจึงช่วย

เพิ่มประสิทธิภาพในการวิเคราะห์ความรู้สึกจากข้อความที่มีความซับซ้อนทางลำดับภาษา เช่น ประโยคที่มีการเปลี่ยนน้ำเสียงหรือมีถ้อยคำแฝง

5. FastText + GRU เป็นการใช้เทคนิคการฝังคำจาก FastText ซึ่งสามารถสร้างเวกเตอร์จากคำย่อยและเข้าใจคำใหม่ ๆ ได้ดี ผสานกับโครงข่าย GRU (Gated Recurrent Unit) ซึ่งเป็นโครงข่ายประสาทเทียมแบบลำดับที่มีโครงสร้างใกล้เคียงกับ LSTM แต่มีจำนวนพารามิเตอร์น้อยกว่า จึงสามารถฝึกโมเดลได้รวดเร็วและใช้ทรัพยากรน้อยลง ขณะเดียวกันยังคงความสามารถในการเรียนรู้บริบทจากลำดับข้อมูลอย่างมีประสิทธิภาพ การใช้ FastText ร่วมกับ GRU จึงเป็นทางเลือกที่เหมาะสมสำหรับการวิเคราะห์ความรู้สึกในกรณีที่ต้องการความสมดุลระหว่างประสิทธิภาพและความเร็วในการประมวลผล โดยยังรองรับข้อความภาษาไทยที่มีลักษณะไม่เป็นทางการได้ดี

3.5.1 โครงสร้างทั่วไปของทุกแบบจำลอง

ค่าพารามิเตอร์เบื้องต้นของแบบจำลองทั้ง 5 เทคนิค ดังตาราง 7

ตาราง 7 ค่าพารามิเตอร์เบื้องต้นของแบบจำลองทั้ง 5 เทคนิค

Parameter	FastText	WangchanBERTa
Batch Size	32	16 (เพื่อลดภาระ GPU)
Epochs	10	4 – 5 (เนื่องจากโมเดลขนาดใหญ่)
Learning Rate	0.001	2e-5 (มาตรฐานของ BERT)
Optimizer	Adam	AdamW (สำหรับ BERT-based)
Dropout	0.5	0.3
Hidden Units (LSTM/GRU)	128	128
Bidirectional	True (เฉพาะ BiLSTM)	True (เฉพาะ BiLSTM)
Max Sequence Length	100	128
Word Embedding Dimension	300 (FastText pre-trained)	768 (WangchanBERTa)
Embedding Type	Static (ไม่ fine-tune)	Contextual (pre-trained BERT)
Loss Function	Categorical Crossentropy	Categorical Crossentropy

Parameter	FastText	WangchanBERTa
Activation Function (Output)	Softmax	Softmax
Validation Split	0.2	0.2
Early Stopping	patience = 3	patience = 2

3.5.2 การปรับค่าพารามิเตอร์ของแบบจำลอง (Fine-tuning Hyperparameters) ของ FastText + LSTM / BiLSTM / GRU

การปรับค่าไฮเปอร์พารามิเตอร์ต่างๆ ของแบบจำลองทั้ง 3 เทคนิค เพื่อหาประสิทธิภาพของแบบจำลองที่ดีที่สุดมีดังตาราง 8

ตาราง 8 Fine-tuning Hyperparameters ของ FastText + LSTM / BiLSTM / GRU

Hyperparameter	ช่วงแนะนำในการปรับ (Search Range)
max_len	128
batch_size	8
epochs	10
validation_split	0.2
early_stop.patience	2
Layer 1	64, 128
Layer 2	32, 64
learning_rate	2e-5, 3e-5
dropout	0.5
optimizer	AdamW

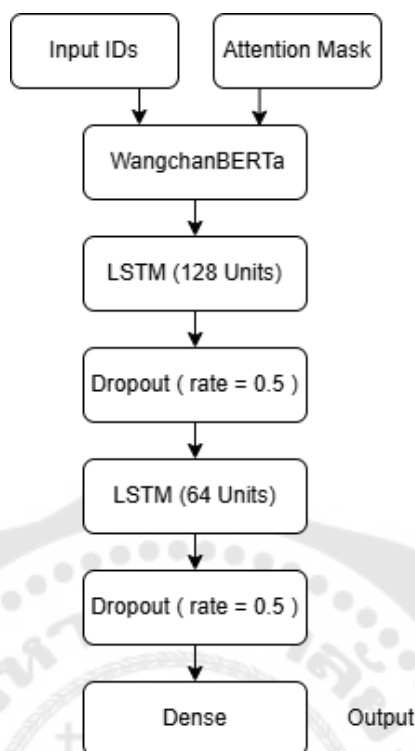
3.5.3 การปรับค่าพารามิเตอร์ของแบบจำลอง (Fine-tuning Hyperparameters) ของ WangchanBERTa + LSTM / BiLSTM

การปรับค่าไฮเปอร์พารามิเตอร์ของแบบจำลองทั้ง 2 เทคนิค เพื่อหาประสิทธิภาพของแบบจำลองที่ดีที่สุดมีดังตาราง 9

ตาราง 9 Fine-tuning Hyperparameters ของ WangchanBERTa + LSTM / BiLSTM

Hyperparameter	ช่วงแนะนำในการปรับ
max_len	128
batch_size	8
epochs	10
validation_split	0.2
early_stop.patience	2
Layer 1	64, 128
Layer 2	32, 64
learning_rate	2e-5, 3e-5
dropout	0.5
optimizer	AdamW
warmup_steps	10% ของ train_steps

สถาปัตยกรรมของแบบจำลองในการปรับค่าไฮเปอร์พารามิเตอร์ สามารถเขียนเป็น Layer ได้ทั้งหมด 8 Layer ดังภาพประกอบ 12



ภาพประกอบ 12 ภาพประกอบสถาปัตยกรรม

3.6 การประเมินผลแบบจำลอง

การประเมินประสิทธิภาพของแบบจำลองการวิเคราะห์ความรู้สึก ผู้วิจัยเลือกใช้ค่าชี้วัดหลัก ได้แก่ Accuracy, Precision, Recall และ F1-score เนื่องจากเป็นตัวชี้วัดมาตรฐานที่สามารถสะท้อนผลลัพธ์ได้รอบด้านและเหมาะสมกับลักษณะของปัญหาการจำแนกข้อความหลายกลุ่ม (Multiclass Classification) โดยค่า Accuracy ให้ความถูกต้องโดยรวมของแบบจำลองในการทำนายข้อความทั้งหมด อย่างไรก็ตาม ในกรณีที่ข้อมูลมีความไม่สมดุลระหว่างกลุ่ม เช่น กลุ่มข้อความเชิงบวกมีจำนวนมากกว่ากลุ่มข้อความเชิงลบหรือกลาง การพิจารณาเพียงค่า Accuracy อาจนำไปสู่การตีความผลลัพธ์ที่คลาดเคลื่อนได้ (Saito & Rehmsmeier, 2015)

เพื่อให้การประเมินผลมีความรอบด้านยิ่งขึ้น ผู้วิจัยจึงพิจารณาค่าชี้วัดเพิ่มเติม ได้แก่ Precision และ Recall ซึ่งเป็นดัชนีที่สามารถสะท้อนมุมมองด้านความแม่นยำของการทำนายคลาสเป้าหมาย และความสามารถในการดึงข้อความจากคลาสนั้นได้ครบถ้วนตามลำดับ โดยเฉพาะในกรณีที่ข้อความมีลักษณะคลุมเครือ หรือแสดงอารมณ์แฝงผ่านภาษาที่ไม่เป็นทางการ การประเมินผลลัพธ์ของโมเดลจึงจำเป็นต้องคำนึงถึงทั้งสองมิติร่วมกัน

นอกจากนี้ เพื่อหาค่าที่แสดงสมดุลระหว่าง Precision และ Recall ได้อย่างเหมาะสม งานวิจัยนี้จึงใช้ F1-score ซึ่งเป็นค่าเฉลี่ยเชิงฮาร์โมนิกของทั้งสองค่านี้ โดยเฉพาะในบริบทที่ความผิดพลาดทั้งประเภทการทำนายเกิน (False Positive) และการทำนายน้อยไป (False Negative) ล้วนมีความสำคัญใกล้เคียงกัน (Sokolova & Lapalme, 2009) วิธีการประเมินด้วยค่าชี้วัดทั้งสองนี้จึงช่วยให้สามารถสะท้อนประสิทธิภาพของแบบจำลองได้ทั้งในระดับภาพรวมและในระดับการวิเคราะห์เชิงลึก

ค่าชี้วัดเหล่านี้ยังได้รับการยอมรับอย่างแพร่หลายในวงวิชาการด้านการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) และการเรียนรู้ของเครื่อง (Machine Learning) โดยเฉพาะในงานที่เกี่ยวข้องกับการจำแนกข้อความ ซึ่งสอดคล้องกับวัตถุประสงค์ของงานวิจัยฉบับนี้ที่มุ่งเน้นการประเมินผลแบบจำลองในการวิเคราะห์ความรู้สึกจากข้อความภาษาไทยที่มีความหลากหลายทั้งในเชิงถ้อยคำและอารมณ์

ดังนั้น งานวิจัยนี้จึงใช้วิธีการประเมินผลแบบจำลองโดยอาศัยค่าชี้วัด ได้แก่ Accuracy, Precision, Recall และ F1-score เพื่อเปรียบเทียบประสิทธิภาพของแต่ละโมเดลและพิจารณาเลือกเทคนิคที่ให้ผลลัพธ์ที่ดีที่สุดสำหรับการจำแนกข้อความในบริบทของข้อมูลภาษาไทยจากรีวิวโรงแรม

บทที่ 4

ผลการศึกษา

การวิเคราะห์ความรู้สึกของข้อความรีวิวโรงแรมบนเกาะภูเก็ตอาศัยกระบวนการเรียนรู้เชิงลึก (Deep Learning) ที่ได้ถูกออกแบบและดำเนินการให้สอดคล้องกับวัตถุประสงค์ของการวิจัยที่มุ่งเน้นการจำแนกความคิดเห็นของผู้ใช้บริการออกเป็นสามระดับ ได้แก่ เชิงบวก เป็นกลาง และเชิงลบ ซึ่งดำเนินการโดยใช้แบบจำลองการเรียนรู้เชิงลึกลำดับจำนวน 3 ประเภท ได้แก่ LSTM (Long Short-Term Memory), BiLSTM (Bidirectional LSTM) และ GRU (Gated Recurrent Unit) ร่วมกับเทคนิคการฝังข้อความ 2 รูปแบบ ได้แก่ FastText และ WangchanBERTa ซึ่งเป็นโมเดลที่ได้รับการฝึกด้วยข้อมูลภาษาไทยโดยเฉพาะ โดยผลลัพธ์ที่ได้จากการวิเคราะห์แบบจำลองร่วมกับเทคนิคการฝังข้อความในแต่ละรูปแบบแบ่งออกเป็น 6 กรณี ได้แก่

1. ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM
 2. ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM
 3. ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM
 4. ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM
 5. ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU
 6. การเปรียบเทียบผลการวิเคราะห์ความรู้สึกระหว่าง LSTM, BiLSTM และ GRU ด้วย FastText และ WangchanBERTa
- โดยมีรายละเอียดผลการวิเคราะห์ในแต่ละกรณี ดังนี้

4.1 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM

แบบจำลองการวิเคราะห์ความรู้สึกจากการตัดคำด้วย WangchanBERTa และใช้โมเดล LSTM โดยจะทำการทดสอบ 2 แบบ คือ

1. การใช้ค่าพารามิเตอร์คงที่

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง WangchanBERTa ร่วมกับ LSTM ภายใต้ค่าพารามิเตอร์คงที่ ดังภาพประกอบ 13 พบว่า มีค่า Accuracy อยู่ที่ 0.599 แสดงถึงระดับความแม่นยำปานกลาง โดยกลุ่มข้อความเชิงบวกให้ผลลัพธ์ที่ดีที่สุด (F1-score = 0.692) รองลงมาคือกลุ่มเชิงลบ (F1-score = 0.627) ขณะที่กลุ่มข้อความเป็นกลางมีความแม่นยำต่ำที่สุด (F1-score = 0.474) ซึ่งอาจเกิดจากลักษณะเนื้อหาที่คลุมเครือหรือเป็นกลาง ทั้งนี้ อาจเกิดจากบริบทของภาษาที่ไม่ชัดเจนในกลุ่มข้อความเป็นกลาง รวมถึงปริมาณข้อมูลในกลุ่มดังกล่าวที่อาจยังไม่เพียงพอ ทั้งนี้ ค่าเฉลี่ยแบบ Macro

และ Weighted เฉลี่ยของ F1-score อยู่ในช่วงประมาณ 0.59 สะท้อนความสมดุลในระดับหนึ่งระหว่างคลาสต่าง ๆ

	Precision	Recall	F1-Score	Support
Negative	0.607	0.648	0.627	193
Neutral	0.491	0.457	0.474	188
Positive	0.692	0.692	0.692	185

Accuracy			0.599	566
Macro avg	0.597	0.599	0.597	566
Weighted avg	0.596	0.599	0.593	566

ภาพประกอบ 13 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM ด้วยค่าคงที่

2. การใช้ Fine-tuning Hyperparameters

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง WangchanBERTa ร่วมกับ LSTM หลังจากการปรับค่าพารามิเตอร์ ดังภาพประกอบ 14 พบว่า มีค่า Accuracy อยู่ที่ 0.694 ซึ่งสูงกว่าค่าที่ได้จากการตั้งค่าพารามิเตอร์แบบคงที่อย่างชัดเจน โดยกลุ่มข้อความเชิงบวกยังคงให้ผลลัพธ์ที่ดีที่สุด (F1-score = 0.798) รองลงมาคือกลุ่มเชิงลบ (F1-score = 0.702) ขณะที่กลุ่มเป็นกลางให้ค่า F1-score เท่ากับ 0.580 แม้จะยังต่ำที่สุดในทั้งสามคลาส แต่มีแนวโน้มดีขึ้นจากค่าพื้นฐานก่อนการ fine-tuning

ค่าเฉลี่ยแบบ Macro และ Weighted ของ Precision, Recall และ F1-score เพิ่มขึ้นมาอยู่ที่ระดับ 0.693 แสดงให้เห็นว่าแบบจำลองมีความสมดุลและมีประสิทธิภาพโดยรวมที่ดีขึ้นหลังจากมีการปรับค่าพารามิเตอร์ ผลลัพธ์นี้สะท้อนว่าแนวทางการ fine-tune แบบจำลองมีความสำคัญต่อการยกระดับความสามารถในการจำแนกความรู้สึกจากข้อความภาษาไทยในบริบทของรีวิวโรงแรมอย่างมีนัยสำคัญ

	Precision	Recall	F1-Score	Support
Negative	0.709	0.694	0.702	193
Neutral	0.591	0.569	0.580	188
Positive	0.776	0.822	0.798	185

Accuracy			0.694	566
Macro avg	0.692	0.695	0.693	566
Weighted avg	0.692	0.694	0.693	566

ภาพประกอบ 14 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + LSTM ด้วยการทำ Fine Tuning Hyperparameter

4.2 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM

แบบจำลองการวิเคราะห์ความรู้สึกจากการตัดคำด้วย WangchanBERTa และใช้โมเดล BiLSTM โดยจะทำการทดสอบ 2 แบบ คือ

1. การใช้ค่าพารามิเตอร์คงที่

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง WangchanBERTa ร่วมกับ BiLSTM ภายใต้ค่าพารามิเตอร์คงที่ ดังภาพประกอบ 15 พบว่ามีค่า Accuracy เท่ากับ 0.572 ซึ่งอยู่ในระดับปานกลาง

ค่าประสิทธิภาพเฉลี่ยแบบ Macro และ Weighted อยู่ที่ประมาณ 0.578 และ 0.577 ตามลำดับ แสดงให้เห็นว่าแบบจำลองยังมีข้อจำกัดในการจำแนกความรู้สึกเมื่อไม่ได้มีการปรับพารามิเตอร์เพิ่มเติม เมื่อพิจารณารายคลาส กลุ่มข้อความเชิงบวกให้ค่า F1-score สูงสุดที่ 0.646 รองลงมาคือกลุ่มเชิงลบ (F1-score = 0.568) ส่วนกลุ่มเป็นกลางยังคงให้ผลลัพธ์ต่ำที่สุด (F1-score = 0.520) แม้จะมีค่า Recall สูงถึง 0.628 ซึ่งบ่งชี้ว่าแบบจำลองสามารถจับข้อความเป็นกลางได้ในระดับหนึ่ง แต่มีความแม่นยำต่ำ ส่งผลให้ F1-score โดยรวมลดลง

	Precision	Recall	F1-Score	Support
Negative	0.681	0.487	0.568	193
Neutral	0.444	0.628	0.520	188
Positive	0.691	0.605	0.646	185

Accuracy			0.572	566
Macro avg	0.605	0.573	0.578	566
Weighted avg	0.606	0.572	0.577	566

ภาพประกอบ 15 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM ด้วยค่าคงที่

2. การใช้ Fine-tuning Hyperparameters

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง WangchanBERTa ร่วมกับ BiLSTM หลังจากการปรับค่าพารามิเตอร์ ดังภาพประกอบ 16 ซึ่งส่งผลให้ค่า Accuracy เพิ่มขึ้นเป็น 0.618 สูงกว่าค่าที่ได้จากการตั้งค่าพารามิเตอร์แบบคงที่ โดยกลุ่มข้อความเชิงบวกให้ผลลัพธ์ดีที่สุด (F1-score = 0.704) รองลงมาคือกลุ่มเชิงลบ (F1-score = 0.685) ส่วนกลุ่มข้อความเป็นกลางยังคงให้ค่า F1-score ต่ำที่สุดที่ 0.377 แม้มีการปรับค่าพารามิเตอร์แล้วก็ตาม

ค่าเฉลี่ยแบบ Macro (F1-score = 0.588) และ Weighted (F1-score = 0.610) ชี้ให้เห็นว่าแบบจำลองมีประสิทธิภาพโดยรวมที่ดีขึ้นเล็กน้อยเมื่อเทียบกับกรณีไม่ปรับค่าพารามิเตอร์ อย่างไรก็ตาม ความแม่นยำในการจำแนกข้อความเป็นกลางยังอยู่ในระดับต่ำ ซึ่งสะท้อนถึงความท้าทายของการแยกแยะข้อความที่มีอารมณ์คลุมเครือ โดยรวมแล้ว ผลลัพธ์นี้แสดงให้เห็นว่าแม้การ Fine-tuning จะช่วยปรับปรุงสมรรถนะของแบบจำลอง BiLSTM ได้ แต่ความแตกต่างระหว่างคลาสยังคงมีอยู่

	Precision	Recall	F1-Score	Support
Negative	0.652	0.721	0.685	193
Neutral	0.420	0.341	0.377	188
Positive	0.695	0.713	0.704	185

Accuracy			0.618	566
Macro avg	0.589	0.592	0.588	566
Weighted avg	0.606	0.618	0.610	566

ภาพประกอบ 16 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย WangchanBERTa + BiLSTM ด้วยการทำ Fine Tuning Hyperparameter

4.3 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM

แบบจำลองการวิเคราะห์ความรู้สึกจากการตัดคำด้วย FastText และใช้โมเดล LSTM โดยจะทำการทดสอบ 2 แบบ คือ

1. การใช้ค่าพารามิเตอร์คงที่

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง FastText ร่วมกับ LSTM โดยใช้ค่าพารามิเตอร์คงที่ ดังภาพประกอบ 17 พบว่าค่า Accuracy เท่ากับ 0.468 จัดอยู่ในระดับต่ำเมื่อเทียบกับแบบจำลองอื่นในชุดการทดลองเดียวกัน โดยเฉพาะค่า F1-score ของข้อความที่เป็นกลางซึ่งต่ำมากเพียง 0.145 แม้ Precision จะอยู่ที่ 0.500 แสดงให้เห็นว่าแบบจำลองมีความแม่นยำในการเลือกคลาสเป็นกลางแต่ไม่สามารถดึงข้อความกลุ่มนี้ได้อย่างครอบคลุม (Recall = 0.085)

สำหรับกลุ่มข้อความเชิงลบและเชิงบวก ค่า F1-score อยู่ที่ 0.544 และ 0.548 ตามลำดับ ซึ่งแม้จะอยู่ในระดับใกล้เคียงกัน แต่ก็ยังสะท้อนประสิทธิภาพโดยรวมที่ไม่สูง ค่าเฉลี่ย Macro F1-score เท่ากับ 0.412 และ Weighted F1-score เท่ากับ 0.413 ซึ่งตอกย้ำข้อจำกัดของแบบจำลองนี้เมื่อไม่ได้รับการปรับแต่งพารามิเตอร์อย่างเหมาะสม โดยสรุป โมเดล FastText + LSTM ที่ใช้ค่าคงที่ยังไม่สามารถจำแนกความรู้สึกจากข้อความภาษาไทยได้อย่างมีประสิทธิภาพ และควรได้รับการปรับปรุงในขั้นตอน Fine-tuning เพื่อเพิ่มศักยภาพในการจำแนกอารมณ์จากข้อมูล

	Precision	Recall	F1-Score	Support
Negative	0.544	0.544	0.544	193
Neutral	0.500	0.085	0.145	188
Positive	0.422	0.778	0.548	185

Accuracy			0.468	566
Macro avg	0.489	0.469	0.412	566
Weighted avg	0.490	0.468	0.413	566

ภาพประกอบ 17 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM ด้วยค่าคงที่

2. การใช้ Fine-tuning Hyperparameters

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง FastText ร่วมกับ LSTM หลังจากการปรับค่าพารามิเตอร์ ดังภาพประกอบ 18 ซึ่งแม้จะไม่ได้ส่งผลกระทบต่อค่า Accuracy โดยรวมที่ยังคงอยู่ที่ 0.468 เช่นเดียวกับกรณีกำหนดค่าพารามิเตอร์คงที่ แต่พบว่าค่าประสิทธิภาพในบางกลุ่มมีการปรับตัวดีขึ้น โดยเฉพาะคลาส Neutral ที่ค่า F1-score เพิ่มขึ้นเป็น 0.527 จากเดิมที่อยู่เพียง 0.145

กลุ่มข้อความเชิงลบให้ค่า F1-score ที่สูงที่สุดคือ 0.600 ขณะที่กลุ่มเชิงบวกแม้จะมี Precision และ Recall สูง (0.718 และ 0.703 ตามลำดับ) แต่ค่า F1-score อยู่ที่ 0.548 ซึ่งแสดงให้เห็นถึงความไม่สมดุลในการทำนายกลุ่มนี้อย่างมีนัยสำคัญ ค่าเฉลี่ย Macro และ Weighted F1-score อยู่ที่ประมาณ 0.412 และ 0.413 ตามลำดับ ซึ่งใกล้เคียงกับกรณีก่อนการ Fine-tuning แม้การปรับค่าพารามิเตอร์จะช่วยเพิ่มความสามารถในการจับข้อความที่เป็นกลาง แต่โดยรวมยังไม่เพียงพอในการยกระดับประสิทธิภาพของแบบจำลอง FastText + LSTM อย่างชัดเจนซึ่งสะท้อนถึงข้อจำกัดของโมเดลเมื่อใช้กับข้อมูลภาษาไทยในบริบทที่มีความซับซ้อนทางอารมณ์

	Precision	Recall	F1-Score	Support
Negative	0.647	0.56	0.600	193
Neutral	0.491	0.569	0.527	188
Positive	0.718	0.703	0.548	185

Accuracy			0.468	566
Macro avg	0.619	0.610	0.412	566
Weighted avg	0.618	0.610	0.413	566

ภาพประกอบ 18 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + LSTM ด้วยการทำ Fine Tuning Hyperparameter

4.4 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM

แบบจำลองการวิเคราะห์ความรู้สึกจากการตัดคำด้วย FastText และใช้โมเดล LSTM โดยจะทำการทดสอบ 2 แบบ คือ

1. การใช้ค่าพารามิเตอร์คงที่

การวิเคราะห์ความรู้สึกของข้อความรีวิวโดยใช้แบบจำลอง FastText ร่วมกับ BiLSTM ภายใต้อัตราการเรียนรู้ที่ ดังภาพประกอบ 19 พบว่ามีค่า Accuracy เท่ากับ 0.473 ซึ่งอยู่ในระดับปานกลาง เมื่อเปรียบเทียบกับแบบจำลองอื่นๆ ที่ใช้ค่าพารามิเตอร์คงที่เช่นกัน โดยกลุ่มข้อความที่ให้ผลดีที่สุดคือ กลุ่มเชิงบวก ซึ่งมี F1-score เท่ากับ 0.517 รองลงมาคือกลุ่มเป็นกลาง (F1-score = 0.501) และกลุ่มเชิงลบให้ผลต่ำที่สุดที่ 0.372 อย่างไรก็ตาม แม้กลุ่มเป็นกลางจะมีค่า Recall สูงถึง 0.649 แต่ Precision อยู่ในระดับต่ำ ส่งผลให้ค่า F1-score ไม่โดดเด่น ค่าเฉลี่ย Macro และ Weighted ของ F1-score อยู่ที่ประมาณ 0.463 และ 0.462 ตามลำดับ ซึ่งสะท้อนว่าแบบจำลองสามารถจำแนกความรู้สึกได้ในระดับพื้นฐานเท่านั้น และยังมีข้อจำกัดในการแยกแยะความรู้สึกที่มีลักษณะคลุมเครือหรือซับซ้อน

	Precision	Recall	F1-Score	Support
Negative	0.576	0.275	0.372	193
Neutral	0.408	0.649	0.501	188
Positive	0.531	0.503	0.517	185

Accuracy			0.473	566
Macro avg	0.505	0.475	0.463	566
Weighted avg	0.506	0.473	0.462	566

ภาพประกอบ 19 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM ด้วยค่าคงที่

2. การใช้ Fine-tuning Hyperparameters

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง FastText ร่วมกับ BiLSTM หลังจากการปรับค่าพารามิเตอร์ ดังภาพประกอบ 20 พบว่าค่า Accuracy อยู่ที่ 0.597 ซึ่งเพิ่มขึ้นอย่างมีนัยสำคัญจากกรณีใช้ค่าพารามิเตอร์คงที่ โดยเฉพาะในกลุ่มข้อความเชิงบวกที่มีค่า F1-score สูงถึง 0.720 พร้อมด้วยค่า Recall ที่ 0.784 ซึ่งแสดงถึงความสามารถของแบบจำลองในการตรวจจับข้อความเชิงบวกได้ดี สำหรับกลุ่มข้อความเชิงลบ ค่า F1-score อยู่ที่ 0.594 ซึ่งสูงกว่าค่าที่ได้จากการตั้งค่าคงที่เช่นกัน ส่วนกลุ่มข้อความเป็นกลาง แม้ค่า F1-score จะยังต่ำที่สุด (0.466) แต่ก็ยังสูงขึ้นจากการทดลองก่อนหน้านี้ แสดงให้เห็นว่าโมเดลมีแนวโน้มเข้าใจข้อความที่คลุมเครือได้ดีขึ้นในระดับหนึ่ง

ค่าประสิทธิภาพโดยรวมจากค่า Macro และ Weighted F1-score เท่ากับ 0.593 สะท้อนว่าแบบจำลอง FastText + BiLSTM หลังการ Fine-tuning มีความสามารถในการจำแนกอารมณ์จากข้อความภาษาไทยได้ดีขึ้นอย่างชัดเจน และมีความสมมูลมากขึ้น เมื่อเทียบกับกรณีก่อนการปรับแต่งพารามิเตอร์

	Precision	Recall	F1-Score	Support
Negative	0.641	0.554	0.594	193
Neutral	0.475	0.457	0.466	188
Positive	0.665	0.784	0.720	185

Accuracy			0.597	566
Macro avg	0.594	0.599	0.593	566
Weighted avg	0.594	0.597	0.593	566

ภาพประกอบ 20 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + BiLSTM ด้วยการทำ Fine Tuning Hyperparameter

4.5 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU

แบบจำลองการวิเคราะห์ความรู้สึกจากการตัดคำด้วย FastText และใช้โมเดล GRU โดยจะทำการทดสอบ 2 แบบ คือ

1. การใช้ค่าพารามิเตอร์คงที่

การวิเคราะห์ความรู้สึกโดยใช้แบบจำลอง FastText ร่วมกับ GRU ภายใต้ค่าพารามิเตอร์ปกติ ดังภาพประกอบ 21 พบว่าแบบจำลองมีประสิทธิภาพต่ำอย่างมีนัยสำคัญ โดยมีค่า Accuracy รวมเพียง 0.34 ซึ่งถือว่าต่ำที่สุดเมื่อเปรียบเทียบกับแบบจำลองอื่นทั้งหมดในชุดการทดลอง ข้อมูลแสดงให้เห็นว่าโมเดลสามารถจำแนกเฉพาะกลุ่มข้อความเชิงลบเท่านั้น โดยมีค่า Recall เท่ากับ 1.00 และ F1-score เท่ากับ 0.51 ขณะที่กลุ่มข้อความที่เป็นกลางและเชิงบวกถูกทำนายผิดทั้งหมด ส่งผลให้ค่า Precision, Recall และ F1-score เท่ากับ 0.00 ในทั้งสองคลาส

ค่าประสิทธิภาพเฉลี่ยทั้ง Macro และ Weighted F1-score อยู่ในระดับต่ำมาก คือ 0.17 และค่า Precision โดยเฉลี่ยไม่เกิน 0.12 แสดงให้เห็นถึงความไม่สมดุลและการล้มเหลวของแบบจำลองในการเรียนรู้ข้อมูลในหลายคลาส โดยเฉพาะกลุ่มข้อความที่มีอารมณ์ชัดเจนหรือคลุมเครือ การตั้งค่าพารามิเตอร์ในระดับปกตินี้จึงไม่เหมาะสมกับโครงสร้างของ GRU เมื่อใช้งานร่วมกับ FastText สำหรับการประมวลผลข้อความภาษาไทย

	Precision	Recall	F1-Score	Support
Negative	0.34	1.00	0.51	193
Neutral	0.00	0.00	0.00	188
Positive	0.00	0.00	0.00	185

Accuracy			0.34	566
Macro avg	0.11	0.33	0.17	566
Weighted avg	0.12	0.34	0.17	566

ภาพประกอบ 21 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU ด้วยค่าปกติ

2. การใช้ Fine-tuning Hyperparameters

การวิเคราะห์ความรู้สึกด้วยแบบจำลอง FastText ร่วมกับ GRU หลังจากการปรับค่าพารามิเตอร์ดังภาพประกอบ 22 พบว่าประสิทธิภาพของแบบจำลองเพิ่มขึ้นอย่างชัดเจนจากเดิม โดยมีค่า Accuracy รวมเท่ากับ 0.565 ทั้งนี้กลุ่มข้อความที่ให้ผลลัพธ์ที่ดีที่สุดคือกลุ่มเชิงบวก ซึ่งมีค่า F1-score เท่ากับ 0.651 รองลงมาคือกลุ่มเชิงลบ (F1-score = 0.589) ขณะที่กลุ่มข้อความที่เป็นกลางยังคงมีค่า F1-score ต่ำที่สุดที่ 0.453 แต่ก็ถือว่าดีขึ้นอย่างมากเมื่อเทียบกับค่าก่อนการ Fine-tuning ที่โมเดลไม่สามารถแยกแยะคลาสได้เลย

ค่าเฉลี่ย Macro F1-score เท่ากับ 0.564 และ Weighted F1-score เท่ากับ 0.564 เช่นกัน แสดงถึงความสมดุลที่ดีขึ้นระหว่างคลาสต่าง ๆ เมื่อเทียบกับการใช้ค่าพารามิเตอร์คงที่ โดยสรุปแบบจำลอง FastText + GRU ที่ผ่านการ Fine-tuning มีศักยภาพในการจำแนกอารมณ์จากข้อความภาษาไทยได้ในระดับที่น่าพอใจ และสามารถนำไปเปรียบเทียบกับแบบจำลองอื่นในบริบทเดียวกันเพื่อคัดเลือกโมเดลที่เหมาะสมที่สุดได้อย่างมีประสิทธิภาพ

	Precision	Recall	F1-Score	Support
Negative	0.603	0.575	0.589	193
Neutral	0.459	0.447	0.453	188
Positive	0.628	0.676	0.651	185

Accuracy			0.565	566
Macro avg	0.563	0.566	0.564	566
Weighted avg	0.563	0.565	0.564	566

ภาพประกอบ 22 ผลลัพธ์การวิเคราะห์ความรู้สึกด้วย FastText + GRU ด้วยการทำ Fine Tuning Hyperparameter

4.6 การเปรียบเทียบผลการวิเคราะห์ความรู้สึกระหว่าง LSTM และ BiLSTM ด้วย FastText และ WangchanBERTa

ก่อนที่จะทำการปรับพารามิเตอร์ จะมาดูในส่วนของผลลัพธ์ก่อนการปรับค่าพารามิเตอร์ในแต่ละแบบจำลอง โดยจะใช้ค่าทั้ง 4 แบบในการประเมินผลลัพธ์

ตาราง 10 ผลลัพธ์ของแต่ละแบบจำลองก่อนการปรับตัวแปร

	Accuracy	Precision	Recall	F1-Score
FastText + LSTM	46.80%	48.90%	46.90%	41.20%
FastText + BiLSTM	47.30%	50.50%	47.50%	46.30%
WangchanBERTa + LSTM	59.90%	59.70%	59.90%	59.70%
WangchanBERTa + BiLSTM	57.20%	60.50%	57.30%	57.80%
FastText + GRU	33.30%	11.10%	33.30%	16.60%

จากตาราง 10 สามารถสรุปผลการเปรียบเทียบประสิทธิภาพของแต่ละโมเดลจากการทดลองในค่าเริ่มต้นพบว่า

1. การใช้ WangchanBERTa เป็นตัวฝังข้อความ (Embedding) มีประสิทธิภาพดีกว่า FastText อย่างมีนัยสำคัญ

2. โมเดลที่ได้ผลดีที่สุด คือ WangchanBERTa ร่วมกับ LSTM โดยให้ค่า Accuracy สูงสุดที่ 59.90%

3. การเปลี่ยนจาก LSTM เป็น BiLSTM ในกรณีของ WangchanBERTa ไม่ได้เพิ่ม Accuracy โดยตรง และในบางตัวชี้วัด เช่น Recall และ F1-Score อาจลดลงเล็กน้อย

4. สำหรับ FastText + GRU ให้ผลลัพธ์ที่ต่ำที่สุดในทุกตัวชี้วัด โดยมี Accuracy เพียง 33.30% และ F1-Score อยู่ที่ 16.60% สะท้อนว่าโมเดลนี้ที่ใช้คำเริ่มต้นไม่สามารถวิเคราะห์ความรู้สึกได้อย่างมีประสิทธิภาพในชุดข้อมูลนี้

ดังนั้น การเลือกใช้โมเดลฝังข้อความที่มีความสามารถในการเข้าใจบริบทเชิงลึก เช่น WangchanBERTa มีผลต่อความแม่นยำของการจำแนกข้อความมากกว่าการปรับเปลี่ยนโครงสร้างของโมเดลลำดับ (Sequential Model) ระหว่าง LSTM และ BiLSTM

ตาราง 11 ผลหลังจากการปรับพารามิเตอร์

	Accuracy	Precision	Recall	F1-Score
FastText + LSTM	60.10%	61.90%	61.00%	61.20%
FastText + BiLSTM	59.70%	59.40%	59.90%	59.30%
WangchanBERTa + LSTM	69.40%	69.20%	69.50%	69.30%
WangchanBERTa + BiLSTM	61.80%	58.90%	59.20%	58.80%
FastText + GRU	56.50%	56.30%	56.60%	56.40%

จากตาราง 11 หลังจากมีการปรับค่าตัวแปรในแต่ละโมเดลแล้วนั้น การวิเคราะห์ความรู้สึกด้วยโมเดล LSTM, BiLSTM และ GRU แบบ 2 ชั้น โดยใช้เทคนิคการฝังข้อความ FastText และ WangchanBERTa พบว่าการเลือกเทคนิคการฝังข้อความมีผลต่อประสิทธิภาพของโมเดลอย่างมีนัยสำคัญ โดย WangchanBERTa ซึ่งเป็นโมเดลภาษาไทยที่มีความสามารถในการเข้าใจบริบทเชิงลึกสามารถเพิ่มค่า Accuracy และ F1-Score ได้สูงกว่าการใช้ FastText อย่างเห็นได้ชัด

นอกจากนี้ จากการเปรียบเทียบระหว่าง LSTM, BiLSTM และ GRU พบว่าการใช้ LSTM ให้ผลลัพธ์ที่ดีกว่า BiLSTM และ GRU ทั้งในกรณีของ FastText และ WangchanBERTa ซึ่งอาจเนื่องมาจาก

ลักษณะของข้อมูลที่ไม่จำเป็นต้องวิเคราะห์ลำดับย้อนกลับมากนัก หรืออาจเกิดจากโครงสร้างของ BiLSTM ที่ทำให้โมเดลมีความซับซ้อนเกินไปสำหรับชุดข้อมูลนี้

โดยสรุปแล้ว การเลือกใช้ WangchanBERTa ร่วมกับ LSTM เป็นแนวทางที่ให้ผลลัพธ์ที่ดีที่สุดในการวิเคราะห์ความรู้สึกจากข้อความภาษาไทยในงานวิจัยนี้ ซึ่งสามารถนำไปประยุกต์ใช้เพื่อพัฒนาระบบวิเคราะห์ความคิดเห็นในระบบจริงต่อไปในอนาคตได้อย่างมีประสิทธิภาพ



บทที่ 5

สรุปผล อภิปรายผล และข้อเสนอแนะ

งานวิจัยนี้ตระหนักถึงบทบาทของข้อมูลรีวิวออนไลน์ที่ผู้ใช้บริการโรงแรมในพื้นที่เกาะภูเก็ต แสดงออกผ่านแพลตฟอร์มสาธารณะ เช่น Google Maps ซึ่งถือเป็นแหล่งข้อมูลที่เปิดโอกาสให้ผู้ใช้งานสามารถแสดงความรู้สึกและประสบการณ์ในการเข้าพักอย่างอิสระ การวิเคราะห์ข้อความเหล่านี้จึงสามารถสะท้อนระดับความพึงพอใจหรือข้อร้องเรียนที่เกิดขึ้นจริงได้อย่างมีนัยสำคัญ อันนำไปสู่แนวคิดในการนำเทคนิคการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) และการวิเคราะห์ความรู้สึก (Sentiment Analysis) มาใช้ในการจำแนกประเภทของความคิดเห็นออกเป็นเชิงบวก เชิงลบ และเป็นกลาง ซึ่งสามารถนำไปใช้ประโยชน์ในการประเมินคุณภาพการให้บริการ และสร้างองค์ความรู้เชิงนโยบายหรือเชิงกลยุทธ์ให้กับผู้ประกอบการในอุตสาหกรรมท่องเที่ยวไทย งานวิจัยนี้จึงมุ่งเปรียบเทียบประสิทธิภาพของแบบจำลองโครงข่ายประสาทเทียมเชิงลำดับ (Recurrent Neural Networks) ที่ได้รับความนิยมในงาน NLP ได้แก่ LSTM, BiLSTM และ GRU โดยออกแบบให้ใช้ชุดข้อมูลเดียวกันที่ได้จากรีวิวภาษาไทยเป็นข้อมูลนำเข้า โดยพิจารณาขอบเขตการวิเคราะห์เฉพาะโรงแรมบนเกาะภูเก็ต เลือกเฉพาะข้อความที่เป็นภาษาไทยและมีเนื้อหาสาระชัดเจน จากนั้นนำข้อมูลเข้าสู่กระบวนการเตรียมข้อมูลเชิงโครงสร้างและข้อความ ชุดข้อมูลที่ผ่านการเตรียมข้อมูลเรียบร้อยแล้วได้ถูกแบ่งเป็นชุดฝึกสอน (training set) และชุดทดสอบ (testing set) ในสัดส่วน 80:20 ก่อนนำไปฝึกด้วยแบบจำลองทั้ง 3 เทคนิคภายใต้เงื่อนไขที่ควบคุมให้เหมือนกัน จากนั้นประเมินประสิทธิภาพของแต่ละโมเดลด้วยค่าชี้วัด ได้แก่ Accuracy, Precision, Recall และ F1-score เพื่อนำผลลัพธ์ที่ได้มาเปรียบเทียบวิเคราะห์ และอภิปรายถึงความเหมาะสมของแต่ละเทคนิคในบริบทของภาษาไทยและเนื้อหาของรีวิวโรงแรมจริง โดยมีรายละเอียดการสรุปผล การอภิปรายผล และข้อเสนอแนะ ดังนี้

5.1 สรุปผล

จากการทดลองเปรียบเทียบประสิทธิภาพของแบบจำลองการวิเคราะห์ความรู้สึกระหว่าง LSTM, BiLSTM และ GRU โดยใช้เทคนิคการฝังข้อความสองรูปแบบ ได้แก่ FastText และ WangchanBERTa สามารถสรุปผลได้ ดังนี้

1. แบบจำลองที่ใช้ WangchanBERTa ร่วมกับ LSTM มีผลลัพธ์ที่ดีที่สุด และยังแสดงให้เห็นถึงแนวโน้มที่สม่ำเสมอในการเพิ่มความแม่นยำและความครอบคลุมของการจำแนกประเภทความรู้สึกเมื่อใช้เทคนิคฝังข้อความเชิงบริบท ซึ่งสามารถจำแนกความรู้สึกจากข้อความได้อย่างแม่นยำ โดยมีค่า

Accuracy เท่ากับ 69.40% Precision เท่ากับ 69.20% Recall เท่ากับ 69.50% และ F1-score เท่ากับ 69.30% แสดงให้เห็นถึงความสอดคล้องของผลลัพธ์ในทุกตัวชี้วัด (ตาราง 11)

2. การใช้เทคนิค WangchanBERTa ซึ่งเป็นโมเดลฝังความหมายของข้อความเชิงบริบทเฉพาะสำหรับภาษาไทย ช่วยเพิ่มประสิทธิภาพของการจำแนกประเภทความรู้สึกได้ดีกว่าการใช้ FastText อย่างชัดเจนในทุกแบบจำลองและทุกตัวชี้วัด

3. เมื่อเปรียบเทียบระหว่างโครงสร้างของแบบจำลองลำดับ พบว่า LSTM มีแนวโน้มให้ผลลัพธ์ที่ดีกว่า BiLSTM และ GRU ทั้งในกรณีที่ใช้ FastText และ WangchanBERTa โดยเฉพาะในด้านค่าความแม่นยำ (Accuracy) และค่าความครอบคลุม (Recall)

4. ผลการทดลองยังชี้ให้เห็นว่า เทคนิคการฝังข้อความ (embedding method) มีอิทธิพลต่อประสิทธิภาพของแบบจำลองมากกว่าโครงสร้างของโมเดลลำดับ โดยการเลือกใช้ WangchanBERTa สามารถยกระดับประสิทธิภาพได้มากกว่าการเปลี่ยนรูปแบบโครงข่ายจาก LSTM เป็น BiLSTM หรือ GRU

5.2 อภิปรายผล

งานวิจัยนี้แสดงให้เห็นถึงบทบาทสำคัญของการเลือกใช้เทคนิคการฝังข้อความ (embedding technique) ต่อประสิทธิภาพของการวิเคราะห์ความรู้สึกจากข้อความภาษาไทย โดยเฉพาะอย่างยิ่งในกรณีที่ใช้โมเดล WangchanBERTa ซึ่งเป็นโมเดลเชิงบริบทที่พัฒนาจากสถาปัตยกรรม Transformer และได้รับการฝึกด้วยข้อมูลภาษาไทยขนาดใหญ่กว่า 16 GB ครอบคลุมแหล่งข้อมูลที่หลากหลาย เช่น วิกิพีเดีย ข่าวออนไลน์ และบทสนทนา WangchanBERTa จึงสามารถเข้าใจลำดับคำและบริบทเฉพาะของภาษาไทยได้ดีซึ่งกว่าโมเดลฝังคำแบบเดิมอย่าง FastText ที่ฝังคำในระดับ subword และไม่สามารถจับความหมายของคำที่เปลี่ยนแปลงตามบริบทได้อย่างแม่นยำ (Devlin et al., 2019; Lowphansirikul et al., 2021)

ผลการทดลองในงานวิจัยนี้แสดงให้เห็นว่า การใช้โมเดล WangchanBERTa ร่วมกับ LSTM สามารถเพิ่มประสิทธิภาพในการจำแนกความรู้สึกของข้อความภาษาไทยได้ดีกว่าการใช้ FastText อย่างมีนัยสำคัญ เมื่อเปรียบเทียบระหว่าง FastText และ WangchanBERTa พบว่า WangchanBERTa ซึ่งเป็นโมเดลที่พัฒนาขึ้นโดยเฉพาะสำหรับภาษาไทย มีความสามารถในการเข้าใจบริบทของภาษาไทยได้ดีกว่า FastText ที่เป็นโมเดลฝังคำแบบดั้งเดิม โดยผลการทดลองแสดงให้เห็นว่า WangchanBERTa ร่วมกับ LSTM ให้ค่าชี้วัดที่สูงกว่าในทุกตัวชี้วัด ได้แก่ Accuracy, Precision, Recall และ F1-score ซึ่งสอดคล้องกับงานวิจัยของ Khamphakdee และ Seresangtakul (2023) โดยพบว่า การใช้โมเดล BERT-based เช่น WangchanBERTa ให้ประสิทธิภาพที่ดีกว่า FastText อย่างมีนัยสำคัญในการจำแนก

ความรู้สึกของข้อความภาษาไทยในโดเมนโรงแรม โดยเฉพาะในข้อมูลที่มีความหลากหลายของความรู้สึกและรูปแบบข้อความ เช่น ข้อความจากโซเชียลมีเดีย หรือรีวิวจากผู้ใช้งาน จากผลการทดลองและงานวิจัยที่เกี่ยวข้อง สามารถสรุปได้ว่า การใช้โมเดลฝังคำเชิงบริบท (Contextualized Embedding Models) เช่น WangchanBERTa ร่วมกับโครงสร้างโมเดลลำดับ (Sequence Models) อย่าง LSTM มีความได้เปรียบในการเรียนรู้การเปลี่ยนแปลงของความหมายคำในบริบทที่แตกต่างกัน

ในส่วนของการสร้างแบบจำลองลำดับ (sequence model structure) การเปรียบเทียบระหว่าง LSTM, BiLSTM และ GRU พบว่า LSTM มักให้ผลลัพธ์ที่ดีกว่าแบบจำลองอื่นๆ อย่างต่อเนื่องในบริบทของชุดข้อมูลนี้ แม้ว่า BiLSTM จะเพิ่มกลไกการเรียนรู้แบบสองทิศทาง ซึ่งในทางทฤษฎีสามารถเพิ่มศักยภาพในการเรียนรู้ความสัมพันธ์แบบย้อนหลังและล่วงหน้าในลำดับคำได้ แต่ผลการทดลองกลับพบว่า BiLSTM ไม่ได้ช่วยเพิ่มประสิทธิภาพอย่างมีนัยสำคัญ และในบางกรณียังมีค่าความแม่นยำต่ำกว่า LSTM แบบทิศทางเดียวเล็กน้อย ซึ่งอาจอธิบายได้จากลักษณะของข้อมูลในชุดรีวิวที่ใช้ในการศึกษา ซึ่งเป็นข้อความสั้น ความซับซ้อนของลำดับไม่สูง และมีการเปลี่ยนแปลงบริบทไม่มาก ทำให้ LSTM แบบทิศทางเดียวเพียงพอในการจับลำดับความสัมพันธ์ของคำได้อย่างมีประสิทธิภาพ โดยสอดคล้องกับการศึกษาของ Liu et al. (2019) ที่ระบุว่า BiLSTM อาจไม่ให้ผลลัพธ์ที่เหนือกว่าในชุดข้อมูลที่มีข้อความสั้นและมีบริบทไม่ซับซ้อน และยังเพิ่มต้นทุนด้านเวลาในการฝึกโมเดลโดยไม่จำเป็น และในส่วนของผลลัพธ์ที่ได้ไม่ได้นักในส่วนหนึ่งมาจากจำนวนคลาส และการที่มีคลาส Neutral นั้นทำให้เกิดข้อผิดพลาดค่อนข้างมาก โดยเป็นการทายถูก 93 ข้อความ ทายเป็นเชิงลบ 60 ข้อความ และเชิงบวก 35 ข้อความ จากทั้งหมด 188 ข้อความที่เป็นคลาส Neutral และจำนวนข้อความที่นำมาใช้ในงานวิจัยนี้ยังคงน้อยเกินไป

ในภาพรวม ผลการวิจัยนี้ชี้ให้เห็นว่า การเลือกใช้เทคนิคการฝังข้อความที่เหมาะสมกับภาษาเป้าหมายมีผลกระทบต่อผลลัพธ์ของแบบจำลองมากกว่าโครงสร้างของโมเดลลำดับเพียงอย่างเดียว โดยเฉพาะในบริบทของภาษาไทยซึ่งมีลักษณะเฉพาะ เช่น ไม่มีการเว้นวรรคระหว่างคำ ใช้คำซ้อน คำแฝง และมีบริบทเชิงวัฒนธรรมสูง การใช้ embedding model ที่เข้าใจบริบทและโครงสร้างภาษาดังกล่าวจึงเป็นสิ่งสำคัญอย่างยิ่งต่อความแม่นยำในการวิเคราะห์ความรู้สึกจากข้อความ

5.3 ข้อเสนอแนะ

งานวิจัยนี้แสดงให้เห็นถึงบทบาทสำคัญของเทคนิคการฝังข้อความเชิงบริบท (contextualized word embedding) และโครงสร้างแบบจำลองลำดับ (sequential model) ในการวิเคราะห์ความรู้สึกจากข้อความภาษาไทย โดยมีข้อเสนอแนะเพื่อเป็นแนวทางสำหรับการพัฒนางานวิจัยในอนาคต ดังนี้

1. การเลือกใช้โมเดลฝังข้อความที่เหมาะสมกับภาษาไทย ควรพิจารณานำโมเดลที่ได้รับการออกแบบมาเฉพาะสำหรับภาษาไทยมาใช้อย่างต่อเนื่อง เช่น WangchanBERTa, Thai2Fit หรือ ThaiBERTa เนื่องจากโมเดลเหล่านี้ได้รับการฝึกจากชุดข้อมูลภาษาไทยโดยเฉพาะ ซึ่งช่วยให้สามารถเข้าใจโครงสร้างภาษาไทยและบริบทที่แฝงอยู่ได้ดีกว่าโมเดลทั่วไป เช่น FastText หรือ Word2Vec ที่ไม่ได้ออกแบบมาเพื่อบริบทของภาษาไทยโดยตรง การใช้โมเดลเฉพาะทางเช่นนี้ยังช่วยลดความคลาดเคลื่อนในการวิเคราะห์ และเพิ่มความแม่นยำในงานด้าน NLP ภาษาไทยในภาพรวม

2. การเปรียบเทียบกับโมเดลฝังข้อความข้ามภาษา ควรมีการศึกษาวิจัยเปรียบเทียบประสิทธิภาพระหว่างโมเดลฝังคำข้ามภาษา (multilingual models) เช่น XLM-RoBERTa (XLM-R), mBERT หรือ LaBSE กับโมเดลเฉพาะภาษาไทย เพื่อวิเคราะห์ข้อได้เปรียบและข้อจำกัดของแต่ละแนวทาง โดยเฉพาะในบริบทที่มีการใช้งานภาษาผสม เช่น รีวิวที่มีการสลับระหว่างภาษาไทยและอังกฤษ หรือข้อความที่มีโครงสร้างภาษาหลากหลาย ซึ่งโมเดลข้ามภาษาที่ได้รับการฝึกจากหลายภาษาอาจสามารถจับบริบทแบบข้ามภาษาได้ดีกว่า

3. การศึกษาครั้งนี้ใช้โครงข่าย LSTM, BiLSTM และ GRU ซึ่งเป็นแบบจำลองลำดับที่มีประสิทธิภาพพื้นฐาน อย่างไรก็ตาม ในงานวิจัยในอนาคตควรสำรวจโครงสร้างแบบจำลองลำดับที่ซับซ้อนยิ่งขึ้น เช่น Transformer Encoder-only models, Transformer Decoder-based หรือ Attention-based LSTM และ Temporal Convolutional Networks (TCNs) เพื่อวิเคราะห์ผลกระทบของโครงสร้างโมเดลที่สามารถเรียนรู้ลำดับข้อมูลในรูปแบบที่มีความลึกหรือมีกลไกความสนใจ (attention mechanism) อย่างละเอียดมากยิ่งขึ้น

4. งานวิจัยนี้ดำเนินการบนชุดข้อมูลข้อความภาษาไทยที่มีความยาวปานกลางและไม่ซับซ้อนมากนัก ในอนาคตควรทดลองเพิ่มขนาดชุดข้อมูลให้มากขึ้น ทั้งในเชิงปริมาณและความหลากหลายของบริบท เช่น ข้อความสนทนาแบบไม่เป็นทางการ บทวิจารณ์จากแพลตฟอร์มต่าง ๆ หรือข้อความที่มีโครงสร้างซ้อนกันหลายระดับ เพื่อศึกษาว่าแบบจำลองลำดับสองทิศทาง เช่น BiLSTM หรือโมเดลแบบ stacked มีความได้เปรียบในแง่ของข้อมูลที่ซับซ้อนมากขึ้นหรือไม่

บรรณานุกรม

- Ayutthaya, T. S. N., & Pasupa, K. (2018). Thai sentiment analysis via bidirectional LSTM-CNN model with embedding vectors and sentic features. 2018 International joint symposium on artificial intelligence and natural language processing (iSAI-NLP),
- Bojanowski, P., Grave, E., Joulin, A., & Mikolov, T. (2017). Enriching Word Vectors with Subword Information. *Transactions of the Association for Computational Linguistics*, 5, 135-146.
https://doi.org/10.1162/tacl_a_00051
- Bowornlertsutee, P., & Paireekreng, W. (2022). The Model of Sentiment Analysis for Classifying the Online Shopping Reviews. *Journal of Engineering and Digital Technology (JEDT)*, 10(1), 71-79.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. North American Chapter of the Association for Computational Linguistics,
- Dey, R., & Salem, F. M. (2017, 6-9 Aug. 2017). Gate-variants of Gated Recurrent Unit (GRU) neural networks. 2017 IEEE 60th International Midwest Symposium on Circuits and Systems (MWSCAS),
- Faqia, Z., Sibaroni, Y., & Prasetyowati, S. (2024). Sentiment Analysis on TikTok App using Long Short-Term Memory (LSTM) with Stochastic Gradient Descent (SGD) Optimization. *JURNAL MEDIA INFORMATIKA BUDIDARMA*, 8, 1292. <https://doi.org/10.30865/mib.v8i3.7699>
- Fu, R., Zhang, Z., & Li, L. (2016, 11-13 Nov. 2016). Using LSTM and GRU neural network methods for traffic flow prediction. 2016 31st Youth Academic Annual Conference of Chinese Association of Automation (YAC),

Google. (2025). *Google map*.

https://www.google.com/maps/place/Koh+Kood/@11.6620417,102.4841151,30845m/data=!3m2!1e3!4b1!4m6!3m5!1s0x31069badd0f9c685:0x6e4c6c3c1cb65b83!8m2!3d11.6680759!4d102.5642261!16s%2Fg%2F121zp6zs?entry=tu&g_ep=EgoyMDI1MDQzMCA4xIKXMDS0ASAFAw%3D%3D

Hochreiter, S., & Schmidhuber, J. (1997). Long Short-Term Memory. *Neural Computation*, 9(8), 1735-1780. <https://doi.org/10.1162/neco.1997.9.8.1735>

Husein, A. M., Livando, N., Andika, A., Chandra, W., & Phan, G. (2023). Sentiment Analysis Of Hotel Reviews On Tripadvisor With LSTM And ELECTRA. *Sinkron: jurnal dan penelitian teknik informatika*, 7(2), 733-740.

Khamphakdee, N., & Seresangtakul, P. (2021). Sentiment analysis for Thai language in hotel domain using machine learning algorithms. *Acta Informatica Pragensia*, 10(2), 155-171.

Leelawat, N., Jariyapongpaiboon, S., Promjun, A., Boonyarak, S., Saengtabtim, K., Laosunthara, A., Yudha, A. K., & Tang, J. (2022). Twitter data sentiment analysis of tourism in Thailand during the COVID-19 pandemic using machine learning. *Heliyon*, 8(10).

Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). RoBERTa: A Robustly Optimized BERT Pretraining Approach. *ArXiv*, abs/1907.11692.

Lowphansirikul, L., Polpanumas, C., Jantrakulchai, N., & Nutanong, S. (2021). WangchanBERTa: Pretraining transformer-based Thai Language Models. *ArXiv*, abs/2101.09635.

Pennington, J., Socher, R., & Manning, C. (2014, October). GloVe: Global Vectors for Word Representation. In A. Moschitti, B. Pang, & W.

- Daelemans, *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)* Doha, Qatar.
- Saito, T., & Rehmsmeier, M. (2015). The Precision-Recall Plot Is More Informative than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets. *PLOS ONE*, 10(3), e0118432.
<https://doi.org/10.1371/journal.pone.0118432>
- Sokolova, M., & Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4), 427-437.
<https://doi.org/https://doi.org/10.1016/j.ipm.2009.03.002>
- Wikipedia. (2024b, Oct 10). *Recurrent neural network*.
https://en.wikipedia.org/wiki/Recurrent_neural_network
- wikipedia. (2025). *Long short-term memory*. https://en.wikipedia.org/wiki/Long_short-term_memory
- Wu, J., Wen, M., Lu, R., Li, B., & Li, J. (2020, 09/01). *The architecture of fastText used for text classification*.
https://www.researchgate.net/publication/340981616_Toward_efficient_and_effective_bullying_detection_in_online_social_network/figures