

MULTI-LABEL SENTIMENT ANALYSIS MODEL FOR THAI LANGUAGE REVIEW OF RESTAURANT

Nontakul PHETRUCHAENG^{1*}, Dussadee PRASERTTITIPONG² and Wijak SRISUJJALERTWAJA³

1 Department of Computer Science, Faculty of Science, Chaing Mai University, Thailand;
nontakul_p@cmu.ac.th (Corresponding Author) (N. P.); dussadee.p@cmu.ac.th (D. P.);
wijak.s@cmu.ac.th (W. S.)

ARTICLE HISTORY

Received: 7 April 2025

Revised: 21 April 2025

Published: 6 May 2025

ABSTRACT

Currently, Sentiment Analysis has been extensively studied, particularly in the context of customer reviews. However, most studies focus on Single-label Sentiment Analysis, which evaluates the overall sentiment of a review as a single entity. In contrast, restaurant reviews often comprise multiple aspects, such as food quality, price, service, and ambience. This study aims to develop a Multi-label Sentiment Analysis Model for Thai restaurant reviews using the Wongnai Review Dataset. The model's performance is evaluated through Traditional Machine Learning approaches, including Logistic Regression, Random Forest, and Support Vector Machine (SVM), in conjunction with text representation techniques such as Bag of Words (BoW) and Term Frequency-Inverse Document Frequency (TF-IDF). The experimental results show that the Random Forest model with TF-IDF achieves superior accuracy of 97.36% in classifying both aspects and sentiments. This study addresses the research gap in multi-aspect sentiment analysis and contributes to the advancement of Thai Natural Language Processing (NLP) applications in the future.

Keywords: Multi-Label Sentiment Analysis, Thai NLP, Machine Learning, Restaurant Review

CITATION INFORMATION: Phetruchaeng, N., Praserttitipong, D., & Srisujjalertwaja, W. (2025). Multi-label Sentiment Analysis Model for Thai Language Review of Restaurant. *Procedia of Multidisciplinary Research*, 3(5), 13

โมเดลวิเคราะห์อารมณ์หลายป้ายสำหรับรีวิวภาษาไทยของร้านอาหาร

นนทกุล เพชรรู้แจ้ง¹, ดุษฎี ประเสริฐธิตินพงษ์² และ วิจักขณ์ ศรีสัจจะเลิศวาจา³

1 สาขาวิทยาการคอมพิวเตอร์ คณะวิทยาศาสตร์ มหาวิทยาลัยเชียงใหม่; nontakul_p@cmu.ac.th
(Corresponding Author) (นนทกุล); dussadee.p@cmu.ac.th (ดุษฎี); wijak.s@cmu.ac.th (วิจักขณ์)

บทคัดย่อ

ปัจจุบันมีการศึกษาด้านการวิเคราะห์อารมณ์อย่างแพร่หลาย โดยเฉพาะในบริบทของรีวิวลูกค้า อย่างไรก็ตาม ส่วนใหญ่มักมุ่งเน้นการวิเคราะห์อารมณ์ป้ายเดียว ซึ่งพิจารณาอารมณ์โดยรวมของรีวิวเพียงแง่มุมเดียว ในขณะที่รีวิวของร้านอาหารนั้น มักประกอบด้วยหลายแง่มุม เช่น รสชาติอาหาร ราคา การให้บริการ และบรรยากาศ งานวิจัยนี้มีเป้าหมายเพื่อพัฒนาโมเดลวิเคราะห์อารมณ์หลายป้ายสำหรับรีวิวภาษาไทยของร้านอาหาร โดยใช้ชุดข้อมูลรีวิวของแพลตฟอร์มวงใน และประเมินผลโมเดลผ่านการเรียนรู้ของเครื่องแบบดั้งเดิม ได้แก่ การถดถอยโลจิสติกส์ การสุ่มป่าไม้ และซัพพอร์ตเวกเตอร์แมชชีน ร่วมกับเทคนิคคลังคำศัพท์ และเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร ผลลัพธ์จากการทดลองแสดงให้เห็นว่าโมเดลการสุ่มป่าไม้ที่ใช้ร่วมกับเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสารนั้นให้ค่าความถูกต้องสูงสุดที่ 97.36% ในการจำแนกทั้งแง่มุมและอารมณ์อย่างแม่นยำ งานวิจัยนี้ช่วยเติมเต็มช่องว่างของการวิเคราะห์อารมณ์ในรีวิวที่มีหลายแง่มุม และสามารถนำไปประยุกต์ใช้กับงานด้านการประมวลผลภาษาธรรมชาติของภาษาไทยในอนาคต

คำสำคัญ: Multi-label Sentiment Analysis, Thai NLP, Machine Learning, Restaurant Review

ข้อมูลการอ้างอิง: นนนทกุล เพชรรู้แจ้ง, ดุษฎี ประเสริฐธิตินพงษ์ และ วิจักขณ์ ศรีสัจจะเลิศวาจา. (2568). โมเดลวิเคราะห์อารมณ์หลายป้ายสำหรับรีวิวภาษาไทยของร้านอาหาร. *Procedia of Multidisciplinary Research*, 3(5), 13

บทนำ

ในโลกยุคดิจิทัลพฤติกรรมของผู้บริโภคมีความเปลี่ยนแปลงไปอย่างมาก โดยเปลี่ยนจากการนั่งรับประทานอาหารที่ร้าน เป็นการสั่งซื้ออาหารผ่านแพลตฟอร์มออนไลน์มากขึ้น โดยตลอดทั้งปี 2567 ที่ผ่านมามีจำนวนผู้บริโภคทั่วโลกที่ซื้อสินค้าอุปโภคบริโภคผ่านทางระบบออนไลน์เป็นจำนวน 5.35 พันล้านคน (Kemp, 2024) ซึ่งการซื้อผ่านระบบออนไลน์ในปัจจุบันนั้นช่วยอำนวยความสะดวกสบายให้แก่ผู้บริโภคมากขึ้น และปัจจัยสำคัญที่มีอิทธิพลต่อการตัดสินใจซื้อมากที่สุดก็คือ รีวิวลูกค้า (Customer Reviews) ที่แสดงบนหน้ารีวิวของแพลตฟอร์มออนไลน์ (Lovert, 2024)

อย่างไรก็ตาม ข้อมูลรีวิวที่มีอยู่เป็นข้อมูลที่ไม่มีโครงสร้าง (Unstructured Data) และมีความท้าทายหลายประการ เช่น 1) มีหลายแง่มุมในประโยคเดียวกัน (Multi-aspect Review) เช่น “อาหารอร่อย แต่ราคาค่อนข้างแพง บริการดี” ซึ่งครอบคลุมแง่มุมอาหาร ราคา และบริการ 2) ความซับซ้อนของภาษาไทยที่มีลักษณะเฉพาะ เช่น การไม่มีการเว้นวรรคระหว่างคำ การใช้คำซ้อน และการใช้คำแสลง เช่น “เค็กร่อยมวักกกแม่” 3) อารมณ์หรือความรู้สึกที่แตกต่างกันในแต่ละแง่มุม เช่น “พนักงานบริการดี แต่อาหารรสชาติไม่อร่อย” และ 4) การจัดการข้อมูลที่ไม่สมดุล (Imbalanced Data) ซึ่งทำให้การจำแนกอารมณ์ซับซ้อนขึ้น

ปัจจุบันมีงานวิจัยเกี่ยวกับการวิเคราะห์อารมณ์ (Sentiment Analysis) สำหรับรีวิวภาษาไทยเกี่ยวกับร้านอาหารเป็นจำนวนมาก แต่ส่วนใหญ่ยังเป็นการวิเคราะห์อารมณ์ป้ายเดียว (Single-label Sentiment Analysis: SLSA) ที่จำแนกอารมณ์ของรีวิวโดยรวมเท่านั้น (Phuttakij et al., 2025) และยังไม่สามารถจำแนกอารมณ์ของแต่ละแง่มุมในรีวิวได้อย่างแม่นยำ ดังนั้นงานวิจัยนี้มุ่งเน้นไปที่การพัฒนาโมเดลวิเคราะห์อารมณ์หลายป้าย (Multi-label Sentiment Analysis: MLSA) ที่สามารถจำแนกแง่มุมและอารมณ์ของรีวิวภาษาไทยเกี่ยวกับร้านอาหารได้พร้อมกัน เพื่อให้ผู้บริโภคสามารถใช้พิจารณาเลือกซื้อสินค้าได้ง่ายขึ้น และเป็นเครื่องมือที่ช่วยให้เจ้าของธุรกิจเข้าใจจุดแข็งและจุดอ่อนของร้านค้าได้ดียิ่งขึ้นด้วย

จากเหตุผลดังกล่าวจึงทำให้ผู้วิจัยเล็งเห็นถึงความสำคัญในการศึกษาเรื่อง “โมเดลวิเคราะห์อารมณ์หลายป้ายสำหรับรีวิวภาษาไทยของร้านอาหาร” โดยมีวัตถุประสงค์การวิจัยเพื่อ 1) เพื่อพัฒนาโมเดลวิเคราะห์อารมณ์หลายป้ายที่สามารถจำแนกแง่มุมและอารมณ์ของรีวิวภาษาไทยได้พร้อมกัน 2) เพื่อเปรียบเทียบประสิทธิภาพของโมเดลการเรียนรู้ของเครื่องแบบดั้งเดิม ได้แก่ การถดถอยโลจิสติกส์ (Logistic Regression: LR) การสุ่มป่าไม้ (Random Forest: RF) และซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) 3) เพื่อศึกษาผลกระทบของเทคนิคคลังคำศัพท์ (Bag of Words: BoW) และเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (Term Frequency-Inverse Document Frequency: TF-IDF) ในการแปลงข้อความเป็นเวกเตอร์ (Text to Vector Conversion) และ 4) เพื่อวิเคราะห์และพัฒนาแบบจำลองการประเมินโมเดลที่เหมาะสมสำหรับรีวิวภาษาไทย

การทบทวนวรรณกรรม

การวิเคราะห์อารมณ์เป็นกระบวนการที่ใช้ในการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) เพื่อจำแนกความคิดเห็นหรืออารมณ์ของข้อความออกเป็นหมวดหมู่ เช่น เชิงบวก (Positive) เชิงลบ (Negative) หรือเป็นกลาง (Neutral) อย่างไรก็ตาม การจำแนกอารมณ์ของรีวิวโดยภาพรวมอาจไม่เพียงพอสำหรับข้อความที่มีหลายแง่มุม ซึ่งพบได้บ่อยในรีวิวร้านอาหารที่ผู้ใช้แสดงความคิดเห็นเกี่ยวกับหลายประเด็น เช่น รสชาติอาหาร ราคา การบริการ และบรรยากาศ การวิเคราะห์อารมณ์หลายป้ายเป็นแนวทางที่ช่วยให้สามารถจำแนกอารมณ์ที่เกี่ยวข้องกับหลายแง่มุมในข้อความเดียวได้ ซึ่งช่วยให้การวิเคราะห์ความคิดเห็นของลูกค้าละเอียดขึ้น งานวิจัยนี้จึงทบทวนวรรณกรรมที่เกี่ยวข้องกับ 4 ประเด็น ได้แก่

การพัฒนาแนวคิดการวิเคราะห์อารมณ์

งานวิจัยด้านการวิเคราะห์อารมณ์นั้นได้พัฒนาไปตามยุคสมัย โดยมีการเปลี่ยนแปลงสำคัญดังนี้

- ในปี 2002 งานวิจัยช่วงแรก ได้มุ่งเน้นที่การวิเคราะห์อารมณ์แบบไบนารี (Binary Sentiment Analysis) ซึ่งจำแนกความคิดเห็นเป็นเพียงสองประเภทคือ Positive หรือ Negative โดยใช้โมเดลซัพพอร์ตเวกเตอร์แมชชีน (SVM) และนาอิวเบย์ (Naïve Bayes) อย่างไรก็ตาม แนวทางนี้ไม่สามารถจับความหลากหลายของอารมณ์ในข้อความที่ซับซ้อนได้ เช่น รีวิวร้านอาหารที่มีหลายมิติ (Pang et al., 2002)

- ในปี 2012 เพื่อแก้ไขข้อจำกัดของแนวทางไบนารี เหล่านักวิจัยได้มีการพัฒนาการวิเคราะห์อารมณ์หลายคลาส (Multi-class Sentiment Analysis) ซึ่งสามารถแบ่งระดับอารมณ์เป็นหลายคลาส เช่น Positive Neutral Negative หรือระดับความพึงพอใจ แต่อย่างไรก็ตาม วิธีนี้ยังไม่สามารถจับอารมณ์ที่เกี่ยวข้องกับหลายแง่มุมในข้อความเดียวได้ (Liu, 2012)

- ในปี 2014 ได้มีการเสนอการวิเคราะห์อารมณ์หลายป้าย ซึ่งสามารถจำแนกข้อความหนึ่งไปยังหลายป้ายพร้อมกัน เช่น รีวิวร้านอาหารที่อาจมีความเห็นเกี่ยวกับ อาหาร ราคา บริการ และบรรยากาศ งานวิจัยนี้จึงสอดคล้องกับลักษณะของรีวิวที่ซับซ้อนโดยเฉพาะในภาษาไทย (Liu & Chen, 2014)

- ในปี 2019 Phatthiyaphaibun และทีมวิจัยได้ศึกษาแนวทางแก้ไขอุปสรรคในการประมวลผลภาษาไทยสำหรับการวิเคราะห์อารมณ์ ซึ่งถือเป็นการวิเคราะห์อารมณ์สำหรับภาษาไทยที่มีนัยสำคัญ โดยเน้นย้ำถึงความท้าทายทางภาษาอันเป็นเอกลักษณ์ที่เกิดจากภาษาไทย เช่น กระบวนการตัดคำ (Tokenization) การแบ่งส่วนคำ (Word Segmentation) และการจัดการรูปร่างลักษณะของคำพูด (Morphology) ที่หลากหลายของข้อความภาษาไทย ซึ่งถือเป็นความท้าทายที่สำคัญมาก (Phatthiyaphaibun et al., 2019)

- ในปี 2021 งานวิจัยล่าสุดได้แสดงให้เห็นว่า WangchanBERTa ซึ่งเป็นโมเดลภาษาที่ถูกฝึกฝนมาล่วงหน้า (Pre-trained Language Models) ที่พัฒนาเฉพาะสำหรับภาษาไทย สามารถเพิ่มความแม่นยำของการวิเคราะห์อารมณ์หลายป้ายได้อย่างมีนัยสำคัญ ถือเป็นผลงานวิจัยที่ดีในการบรรลุความถูกต้องแม่นยำที่สูงขึ้นและความเข้าใจอารมณ์ของผู้ใช้อย่างละเอียดถี่ถ้วน (Lowphansirikul et al., 2021)

การสกัดคุณลักษณะ (Feature Extraction)

สามารถแบ่งได้เป็น 2 ขั้นตอน ได้แก่

1) กระบวนการตัดคำ ภาษาไทยไม่มีการเว้นวรรคระหว่างคำและมีคำที่มีหลายความหมาย ซึ่งทำให้กระบวนการตัดคำมีความซับซ้อนกว่าภาษาอื่น โดยงานวิจัยนี้ได้ใช้ไลบรารี PyThaiNLP ซึ่งเป็นเครื่องมือที่ถูกออกแบบมาเพื่อรองรับการประมวลผลข้อความภาษาไทยโดยเฉพาะ ประกอบกับงานวิจัยของ Phatthiyaphaibun พบว่า PyThaiNLP มีประสิทธิภาพสูงสุดในการแปลงข้อความภาษาไทย เมื่อเทียบกับ LexTo และ DeepCut (Phatthiyaphaibun et al., 2023)

2) เทคนิคการแปลงข้อความ โดยแนวทางหลักในการแปลงข้อความเป็นเวกเตอร์มี 3 เทคนิค ได้แก่

2.1) เทคนิคคลังคำศัพท์ (BoW) เป็นเทคนิคพื้นฐานในการแปลงข้อความเป็นเวกเตอร์ โดยนับจำนวนการปรากฏของแต่ละคำในเอกสารโดยไม่คำนึงถึงลำดับของคำ ข้อความจะถูกแปลงเป็นเมทริกซ์ (Matrix) ของความถี่ของคำ ซึ่งสามารถใช้เป็นอินพุต (Input) ให้กับโมเดลการเรียนรู้ของเครื่อง (Machine Learning) ได้ แม้ว่าเทคนิคคลังคำศัพท์จะเรียบง่ายและมีประสิทธิภาพ แต่ข้อจำกัดหลักคือการละเลยลำดับของคำและความสัมพันธ์ทางไวยากรณ์ (Sangsavate et al., 2023)

2.2) เทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) เป็นเทคนิคที่พัฒนาต่อยอดจากเทคนิคคลังคำศัพท์ โดยให้ค่าน้ำหนักแก่คำที่มีความสำคัญมากขึ้น คำนวณจากความถี่ของคำในเอกสารหนึ่ง (TF) เทียบกับส่วนกลับความถี่ของคำในเอกสารทั้งหมด (IDF) วิธีนี้ช่วยลดอิทธิพลของคำที่พบบ่อยในทุกเอกสารและเพิ่มความสำคัญให้กับคำที่มีลักษณะเฉพาะของเอกสารนั้นๆ ทำให้สามารถแยกแยะความแตกต่างของเอกสารได้ดีขึ้น (Khamphakdee & Seresangtakul, 2021) ซึ่งสามารถอธิบายเป็นสมการ ได้ดังนี้

a) Term Frequency (TF): คำนวณความถี่ของคำที่ปรากฏในเอกสาร

สมการที่ 1 คำนวณหาค่าความถี่ของคำในเอกสารหนึ่ง

$$TF = \frac{\text{จำนวนครั้งที่คำปรากฏในเอกสาร}}{\text{จำนวนคำทั้งหมดในเอกสาร}}$$

b) Inverse Document Frequency (IDF): คำนวณว่าคำดังกล่าวพบได้ทั่วไปในทุกเอกสารหรือไม่

สมการที่ 2 คำนวณหาค่าส่วนกลับความถี่ของคำในเอกสารทั้งหมด

$$IDF = \log\left(\frac{\text{จำนวนเอกสารทั้งหมด}}{\text{จำนวนเอกสารที่มีคำนี้ปรากฏ}}\right)$$

c) TF-IDF Calculation:

สมการที่ 3 คำนวณหาค่าความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร

$$TF - IDF = TF \times IDF$$

2.3) การฝังคำ (Word Embeddings) เป็นเทคนิคการแปลงคำให้อยู่ในรูปแบบเวกเตอร์ที่สามารถจับความสัมพันธ์ทางความหมายของคำได้ โดยโมเดลเรียนรู้จากบริบทของคำที่ปรากฏในข้อมูลขนาดใหญ่ วิธีนี้ช่วยให้คำที่มีความหมายใกล้เคียงกันมีเวกเตอร์ที่อยู่ใกล้กันในพื้นที่มิติสูง ตัวอย่างเครื่องมือของการฝังคำที่ได้รับความนิยม ได้แก่ Word2Vec GloVe FastText และสำหรับภาษาไทย ได้แก่ Thai2Vec และ WangchanBERTa เทคนิคนี้มีความสามารถในการรักษาความสัมพันธ์ทางความหมาย ทำให้มีประสิทธิภาพสูงกว่าเทคนิคคลังคำศัพท์ และเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสารในการวิเคราะห์ภาษาธรรมชาติ (Marukatat, 2020)

อย่างไรก็ตาม ข้อความรีวิวยาษาไทยมีลักษณะเป็นภาษาพูดเป็นส่วนใหญ่ ซึ่งอาจไม่อยู่ในฐานข้อมูลการฝึกของการฝังคำ ในงานวิจัยนี้จึงมุ่งเน้นไปที่เทคนิคคลังคำศัพท์ และเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสารเป็นหลัก

3) เทคนิคการจำแนกประเภทข้อความ (Classification Methods) ในงานวิจัยนี้ ผู้วิจัยใช้วิธีการเรียนรู้ของเครื่องแบบดั้งเดิม โดยเทคนิคที่ใช้สำหรับโมเดลวิเคราะห์อารมณ์หลายป้าย ประกอบด้วย 3 โมเดล ได้แก่

3.1) การถดถอยโลจิสติกส์ เป็นอัลกอริทึมการจำแนกประเภท (Classification Algorithm) ที่ใช้ในการพยากรณ์ความน่าจะเป็นของข้อมูลที่อยู่ในหมวดหมู่หนึ่งจากสองหมวดหมู่ (Binary Classification) หรือหลายหมวดหมู่ (Multiclass Classification) โดยอาศัยสมการถดถอยเชิงเส้น แต่ใช้ฟังก์ชันโลจิสติก (Logistic Function) แปลงค่าให้อยู่ในช่วง 0 กับ 1 ฟังก์ชันนี้ช่วยให้สามารถตีความผลลัพธ์เป็นค่าความน่าจะเป็นได้ และโมเดลการถดถอยโลจิสติกส์ยังใช้ฟังก์ชันการประมาณค่าความน่าจะเป็นสูงสุด (Maximum Likelihood Estimation: MLE) ในการประมาณค่าสัมประสิทธิ์ของตัวแปรอิสระเพื่อลดข้อผิดพลาดและเพิ่มความแม่นยำของโมเดล ข้อดีของโมเดลนี้คือ มีความง่ายต่อการตีความ คำนวณรวดเร็ว และสามารถใช้เป็นโมเดลฐาน (Baseline Model) ได้ดี ส่วนข้อเสียคือ มีข้อจำกัดในการจำแนกข้อมูลที่ไม่มีเส้นตรงและอาจไวต่อค่าออกนอกเกณฑ์ (Outliers)

3.2) การสุ่มป่าไม้ เป็นโมเดลการเรียนรู้แบบกลุ่ม (Ensemble Learning) ซึ่งใช้การรวมโมเดลหลายตัว (หลายต้นไม้ตัดสินใจ: Decision Trees) เพื่อเพิ่มประสิทธิภาพของการจำแนกประเภทและลดข้อผิดพลาดที่เกิดจากโมเดลต้นไม้ตัดสินใจเดี่ยว (Single Decision Tree) โมเดลนี้สร้างต้นไม้หลายต้นจากการสุ่มตัวอย่างข้อมูลแบบบูตสตรapping (Bootstrapping) และใช้บูตสตรapping เพื่อลดความแปรปรวน (Variance) ของโมเดล ข้อดีของโมเดลนี้คือ ทนต่อปัญหาโอเวอร์ฟิตติง (Overfitting) ได้ดี ทำงานได้ดีกับข้อมูลที่มีความซับซ้อนและมีหลายมิติ (High-dimensional Data) ส่วนข้อเสียคือ ใช้ทรัพยากรในการคำนวณสูงกว่าต้นไม้ตัดสินใจเดี่ยวและอาจไม่สามารถตีความผลลัพธ์ได้ง่ายเท่ากับการถดถอยโลจิสติกส์ (Thiengburanathum & Charoenkwan, 2023)

3.3) ซัพพอร์ตเวกเตอร์แมชชีน เป็นโมเดลการจำแนกประเภทที่อาศัยแนวคิดของระนาบเกิน (Hyperplane) หรือขอบเขตการตัดสินใจ (Decision Boundary) ในการแบ่งข้อมูลออกเป็นกลุ่มต่างๆ โดยเลือกเส้นแบ่งที่สามารถแยกข้อมูลออกจากกันได้ดีที่สุด ซึ่งเรียกว่าการจำแนกขอบเขตสูงสุด (Maximum Margin Classifier) โมเดลซัพพอร์ตเวกเตอร์แมชชีนสามารถทำงานได้ทั้งในปัญหาการจำแนกเชิงเส้น (Linear Classification) และไม่เป็นเชิงเส้น (Non-linear Classification) โดยอาศัยฟังก์ชันเคอร์เนล (Kernel Function) ในการเปลี่ยนข้อมูลให้อยู่ในมิติที่สูงขึ้นเพื่อให้

สามารถแบ่งข้อมูลได้ดีขึ้น ข้อดีของโมเดลนี้คือ มีประสิทธิภาพสูงเมื่อตัวแปรมีจำนวนมาก (High-dimensional Data) ทนต่อปัญหาโอเวอร์ฟิตติ้งได้ดี ส่วนข้อเสียคือ ใช้เวลาฝึกโมเดลนานเมื่อตัวอย่างข้อมูลมีขนาดใหญ่และต้องเลือกฟังก์ชันเคอร์เนล (Kernel Function) ที่เหมาะสมเพื่อให้ทำงานได้ดี

4) การประเมินผลโมเดล (Model Evaluation Metrics) ประกอบด้วย 4 ตัวชี้วัดสำคัญในการประเมินโมเดลวิเคราะห์อารมณ์หลายป้าย (Tao & Fang, 2020) ได้แก่

4.1) ค่าความถูกต้อง (Accuracy) เป็นเมตริกพื้นฐานที่ใช้วัดความถูกต้องของโมเดล โดยคำนวณจากสัดส่วนของจำนวนตัวอย่างที่โมเดลพยากรณ์ถูกต้องทั้งหมดต่อจำนวนตัวอย่างทั้งหมดในชุดข้อมูลทดสอบ ดังที่แสดงในสมการที่ 4

สมการที่ 4 คำนวณวัดความถูกต้องของโมเดล

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

a) True Positive (TP) คือ จำนวนตัวอย่างที่เป็นบวกและโมเดลทำนายถูกต้อง

b) True Negative (TN) คือ จำนวนตัวอย่างที่เป็นลบและโมเดลทำนายถูกต้อง

c) False Positive (FP) คือ จำนวนตัวอย่างที่เป็นลบแต่โมเดลทำนายผิดว่าเป็นบวก

d) False Negative (FN) คือ จำนวนตัวอย่างที่เป็นบวกแต่โมเดลทำนายผิดว่าเป็นลบ

4.2) ค่าความแม่นยำ (Precision) เป็นเมตริกที่ใช้วัดความแม่นยำของโมเดลในกรณีที่ทำนายว่าเป็นกลุ่มบวก (Positive Class) โดยพิจารณาว่าในจำนวนที่โมเดลทำนายเป็นบวก มีจำนวนที่ถูกต้องจริงกี่เปอร์เซ็นต์

สมการที่ 5 คำนวณวัดความแม่นยำของโมเดล

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

4.3) ค่าความไว (Recall หรือ Sensitivity) เป็นเมตริกที่ใช้วัดความสามารถของโมเดลในการตรวจจับตัวอย่างที่เป็นบวก โดยคำนวณจากสัดส่วนของตัวอย่างที่เป็นบวกจริงที่ถูกโมเดลทำนายถูกต้องเมื่อเทียบกับจำนวนบวกทั้งหมดในชุดข้อมูล

สมการที่ 6 คำนวณวัดความครบถ้วนของโมเดล

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

4.4) ค่า F1 (F1-Score) เป็นค่าที่รวม Precision และ Recall เข้าด้วยกันโดยใช้ค่าเฉลี่ยถ่วงน้ำหนักแบบฮาร์มอนิก (Harmonic Mean) เพื่อหาจุดสมดุลระหว่างการลด False Positive และ False Negative

สมการที่ 7 คำนวณวัดความสมดุลระหว่าง ค่าความแม่นยำ และค่าความไว

$$\text{F1 Score} = 2 \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

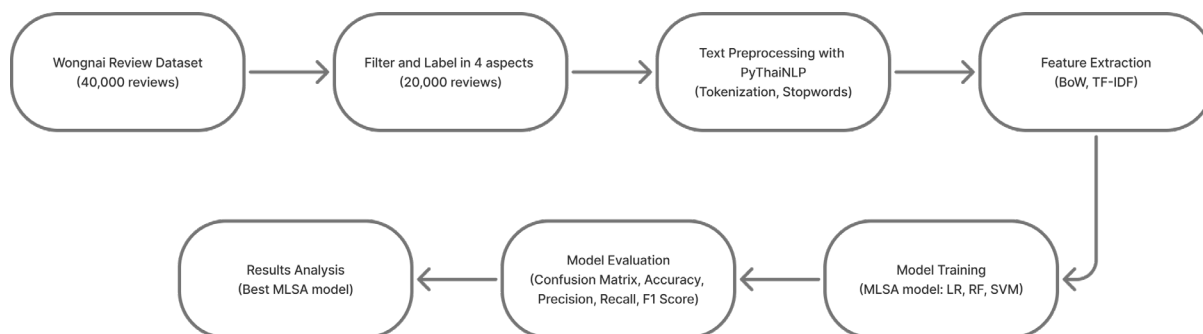
สมมติฐานการวิจัย

1) การใช้โมเดลวิเคราะห์อารมณ์หลายป้าย (MLSA) สามารถเพิ่มประสิทธิภาพในการวิเคราะห์รีวิวภาษาไทยของร้านอาหาร เมื่อเทียบกับการวิเคราะห์อารมณ์ป้ายเดียว (SLSA)

2) เทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) จะช่วยให้โมเดลสามารถจับความหมายของข้อความได้ดีกว่าเทคนิคคลังคำศัพท์ (BoW)

3) โมเดลการสุ่มป่าไม้ (RF) มีแนวโน้มให้ผลลัพธ์ที่แม่นยำกว่าโมเดลการถดถอยโลจิสติกส์ (LR) และโมเดลซัพพอร์ตเวกเตอร์แมชชีน (SVM) สำหรับการวิเคราะห์อารมณ์หลายป้าย

กรอบแนวคิดการวิจัย



ภาพที่ 1 กรอบแนวคิดการวิจัย

วิธีดำเนินการวิจัย

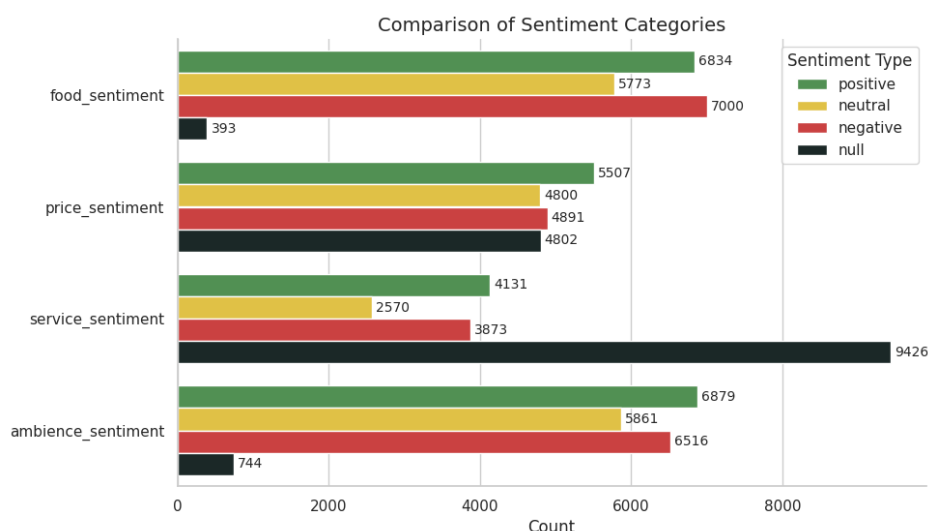
โมเดลวิเคราะห์อารมณ์หลายป้ายสำหรับรีวิวภาษาไทยของร้านอาหารนี้ ประกอบด้วย 6 ขั้นตอน ดังนี้

การเก็บข้อมูล (Data Collection) งานวิจัยนี้แบ่งแ่งมุ่มเป็น 4 ด้าน ได้แก่ อาหาร ราคา บริการ และบรรยากาศ แบ่งอารมณ์เป็น 4 อารมณ์ ได้แก่ positive neutral negative และ null ใช้ชุดข้อมูลรีวิวของแพลตฟอร์มในที่มีรีวิวร้านอาหารจำนวน 40,000 รายการ จากนั้นทำการกำหนดค่าอ็อบเจกต์ (Object value) 4 รูปแบบ ดังนี้ 1) positive 2) neutral 3) negative และ 4) null ลงไปในแต่ละแ่งมุ่มที่นิยมใช้ในงานวิจัยรีวิวร้านอาหาร (Suciati & Budi, 2019) และสุ่มเลือกชุดข้อมูลให้เหลือ 20,000 รายการ ดังตัวอย่างในตารางที่ 1

ตารางที่ 1 ตัวอย่างการกำหนดค่าอ็อบเจกต์ และสุ่มเลือกชุดข้อมูล

id	review	food_sentiment	price_sentiment	service_sentiment	ambience_sentiment
9379	เป็นร้านที่กว้างมาก ๆ มีที่ จอดรถเยอะเลยคะ อาหาร รสชาติดีอร่อยคะ พนักงาน บริการดีมาก ๆ คะ	positive	null	positive	positive
6770	ราคา 351.- อาหาร คุณภาพอาหารต่ำ บริการ ไม่ประทับใจ	negative	neutral	negative	null

เมื่อทำการกำหนดค่าอ็อบเจกต์เรียบร้อยแล้ว จะสามารถสรุปจำนวนของประเภทอารมณ์ (Sentiment Type) ของแต่ละแ่งมุ่ม ได้ดังภาพที่ 2



ภาพที่ 2 เปรียบเทียบจำนวนประเภทอารมณ์ของแต่ละแง่มุม

จากภาพที่ 2 จะเห็นว่าคอลัมน์ service_sentiment มีจำนวน null สูงสุด และตามด้วยคอลัมน์ price_sentiment ซึ่งอาจบ่งชี้ได้หลายกรณี เช่น ข้อมูลบริการไม่ได้ถูกจัดเก็บ หรือ ไม่สามารถให้ความเห็นได้ เป็นต้น ดังนั้นเพื่อจัดการข้อมูลสูญหาย (Missing data) ส่วนนี้จึงทำการแทนค่าด้วย "null" เพื่อให้โมเดลสามารถเรียนรู้ข้อมูลได้ครบจำนวน 20,000 รายการ

การทำความสะอาดข้อมูล (Data Cleaning) ในขั้นตอนนี้ได้ใช้ไลบรารี PyThaiNLP สำหรับกระบวนการตัดคำเพราะเป็นไลบรารีที่ออกแบบมาเพื่อแก้ไขปัญหาของภาษาไทยโดยเฉพาะ พร้อมทั้งทำการลบสัญลักษณ์ที่ไม่จำเป็นและลบคำไร้ประโยชน์ (Stopwords) ที่ไม่มีความหมายในการวิเคราะห์ออกไป

การแปลงข้อความเป็นเวกเตอร์ (Text to Vector Conversion)

แปลงข้อความที่ผ่านการประมวลผลให้เป็นเวกเตอร์เชิงตัวเลขเพื่อให้โมเดลสามารถเรียนรู้ได้ โดยแบ่งเป็น 2 เทคนิค ได้แก่ 1) เทคนิคคลังคำศัพท์ (BoW) ใช้ CountVectorizer เพื่อแปลงข้อความเป็นเมทริกซ์ของการนับจำนวนครั้งของแต่ละคำ และ 2) เทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) ใช้ TfidfVectorizer เพื่อคำนวณค่าน้ำหนักของคำโดยพิจารณาความสำคัญของคำในเอกสาร

การพัฒนาโมเดลวิเคราะห์อารมณ์หลายป้ายนี้ใช้การเรียนรู้ของเครื่องแบบดั้งเดิม ได้แก่ 1) การถดถอยโลจิสติกส์ (LR) 2) การสุ่มป่าไม้ (RF) และ 3) ซัพพอร์ตเวกเตอร์แมชชีน (SVM) ร่วมกับชุดข้อมูลที่ได้จากเทคนิคคลังคำศัพท์ และเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร เพื่อแปลงข้อความเป็นเวกเตอร์ก่อนนำเข้าสู่โมเดล แต่เนื่องจากโมเดลเหล่านี้ไม่ได้ออกแบบมาให้รองรับปัญหาหลายป้ายโดยตรง จึงใช้ไลบรารี MultiOutputClassifier จาก Scikit-learn เพื่อแปลงแต่ละโมเดลให้สามารถรองรับการจำแนกประเภทหลายป้าย (Puttipornchai et al., 2022) โดยไลบรารี MultiOutputClassifier จะทำการฝึกโมเดลเดียวกันสำหรับแต่ละป้ายแยกจากกันอย่างอิสระ ซึ่งทำให้โมเดลนี้สามารถทำนายป้ายและอารมณ์ได้หลายประเภทพร้อมกัน จากนั้นแบ่งชุดข้อมูลสำหรับฝึกฝน (Train) และทดสอบ (Test) เป็นอัตราส่วน 80:20 ฉะนั้นโมเดลวิเคราะห์อารมณ์หลายป้ายนี้จึงมี 6 เทคนิค ได้แก่ 1) LR + BoW 2) LR + TF-IDF 3) RF + BoW 4) RF + TF-IDF 5) SVM + BoW และ 6) SVM + TF-IDF

การประเมินผล (Model Evaluation) หลังจากฝึกโมเดลเรียบร้อยแล้ว จะทำการวัดผลลัพธ์เพื่อประเมินประสิทธิภาพของโมเดล โดยงานวิจัยนี้ได้เลือกตัวชี้วัดที่นิยมใช้กันในการทำโมเดลวิเคราะห์อารมณ์หลายป้าย ซึ่งประกอบด้วย 1) ค่าความถูกต้อง (Accuracy) 2) ค่าความแม่นยำ (Precision) 3) ค่าความไว (Recall) และ 4) ค่า F1 (F1-Score) ดังที่แสดงรายละเอียดในสมการที่ 4-7 ตามลำดับ

การวิเคราะห์เปรียบเทียบผลลัพธ์และเลือกโมเดลที่ดีที่สุด โดยเลือกโมเดลที่มีค่าความถูกต้อง ค่าความแม่นยำ ค่าความไว และค่า F1 สูงสุด โดยสามารถดูสรุปวิธีดำเนินการวิจัยทั้งหมด จากตารางที่ 2

ตารางที่ 2 สรุปวิธีดำเนินการวิจัยโมเดลวิเคราะห์อารมณ์หลายป้ายสำหรับรีวิวภาษาไทยของร้านอาหาร

Algorithm: Multi-label Sentiment Analysis for Thai Restaurant Reviews

INPUT: A Thai restaurant review dataset

OUTPUT: Multi-label sentiment prediction

1 LOAD dataset (Wongnai Review Dataset)

2 REMOVE English characters, special symbols, and numbers

3 REMOVE Thai common words and stopwords

4 SEGMENT text into sentences

5 foreach sentence do

6 TOKENIZE sentence using PyThaiNLP

7 foreach token in sentence do

8 FIND Aspect Dimension [Food, Price, Service, Ambience]

9 if (Aspect is found) then

10 FIND Sentiment of such Aspect [Positive, Neutral, Negative]

11 ADD (Aspect, Sentiment) pair to Preprocessed Data

12 end

13 end

14 end

15 EXTRACT Features using: BoW and TF-IDF

16 DEFINE MultiOutputClassifier with models: LR, RF and SVM

17 TRAIN Multi-label Sentiment Analysis Model on Training Data

18 EVALUATE models using: Accuracy, Precision, Recall and F1-score

19 COMPARE performance of LR, RF, and SVM models

20 SELECT best-performing model based on Accuracy, Precision, Recall, and F1-score

21 END

ผลการวิจัย

จากผลการทดลอง ผูกโมเดลเป็นจำนวน 100 รอบ รวมกับค่าควบคุมข้อมูลสุ่มคือ 42 พบว่า โมเดลที่ใช้เทคนิคการเรียนรู้แบบสุ่มป่าไม้ (Random Forest: RF) รวมกับ การแปลงข้อความด้วยเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) ให้ประสิทธิภาพสูงสุด โดยมีค่าความถูกต้อง (Accuracy) สูงถึง 97.36% ค่าความแม่นยำ (Precision) 97.42% ค่าความไว (Recall) 97.36% และค่าคะแนน F1 (F1-score) 97.34% ซึ่งถือเป็นผลลัพธ์ที่ดีที่สุดในการทดสอบทั้งหมด รองลงมาเป็นคือโมเดลที่ใช้การสุ่มป่าไม้ รวมกับเทคนิคคลังคำศัพท์ (BoW) และลำดับที่สามคือโมเดลที่ใช้ซัพพอร์ตเวกเตอร์แมชชีน (SVM) รวมกับเทคนิคคลังคำศัพท์ (BoW) โดยสามารถสรุปเปรียบเทียบผลลัพธ์เชิงปริมาณของแต่ละโมเดลได้ดังที่แสดงในตารางที่ 3

ในส่วนของการทดสอบการใช้เทคนิคการแปลงข้อความทั้งสองรูปแบบ (TF-IDF และ BoW) ร่วมกัน พบว่าไม่สามารถเพิ่มประสิทธิภาพของโมเดลได้ กล่าวคือ ค่าความถูกต้อง (Accuracy) มีแนวโน้มลดลง และค่าความสูญเสีย (Loss) มี

แนวโน้มเพิ่มขึ้นในแต่ละรอบของการฝึกฝน (Training epoch) ซึ่งสะท้อนถึงปัญหาในการเรียนรู้ของโมเดลในการแปลงข้อความเป็นเวกเตอร์ลักษณะที่เหมาะสม การรวมกันของเทคนิคทั้งสองส่งผลให้โมเดลเกิดความซับซ้อนโดยไม่จำเป็น และลดทอนประสิทธิภาพในการเรียนรู้ลักษณะเฉพาะของข้อมูลข้อความ ดังนั้น ผู้วิจัยจึงเลือกดำเนินการทดสอบโดยใช้เทคนิคการแปลงข้อความแต่ละแบบแยกจากกัน ซึ่งช่วยให้โมเดลสามารถเรียนรู้ได้อย่างมีประสิทธิภาพมากยิ่งขึ้น และส่งผลให้ได้ค่าประเมินที่ดีกว่าอย่างชัดเจน

ตารางที่ 3 เปรียบเทียบผลลัพธ์ของโมเดลวิเคราะห์อารมณ์หลายป้าย

Vectorizer	Model	Accuracy	Precision	Recall	F1-Score
BoW	LR	0.9701	0.9708	0.9701	0.9699
TF-IDF	LR	0.9494	0.9511	0.9494	0.9486
BoW	RF	0.9722	0.9728	0.9722	0.9720
TF-IDF	RF	0.9736	0.9742	0.9736	0.9734
BoW	SVM	0.9718	0.9724	0.9718	0.9717
TF-IDF	SVM	0.9616	0.9626	0.9636	0.9634

สรุปและอภิปรายผลการวิจัย

จากผลการทดลองพบว่าโมเดลวิเคราะห์อารมณ์หลายป้าย (Multi-Label Sentiment Analysis: MLSA) มีประสิทธิภาพสูงกว่าโมเดลวิเคราะห์อารมณ์ป้ายเดียว (Single-Label Sentiment Analysis: SLSA) อย่างมีนัยสำคัญ โดยเฉพาะในบริบทของรีวิวร้านอาหารที่มักประกอบด้วยหลายแง่มุมและสะท้อนอารมณ์หลากหลายภายในข้อความเดียว ซึ่งการวิเคราะห์อารมณ์ป้ายเดียวส่งผลให้เกิดการสูญเสียข้อมูลและไม่สามารถสะท้อนความคิดเห็นที่ซับซ้อนได้อย่างครบถ้วน ในทางตรงกันข้ามการวิเคราะห์อารมณ์หลายป้ายสามารถแยกแยะและระบุอารมณ์ในแต่ละแง่มุมได้อย่างชัดเจนและครอบคลุมมากกว่า

โมเดลที่ให้ผลลัพธ์ที่ดีที่สุดในงานวิจัยนี้คือ โมเดลที่ใช้การสุ่มป่าไม้ (Random Forest: RF) ร่วมกับเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) ซึ่งให้ค่าประเมินประสิทธิภาพสูงในทุกมิติ ได้แก่ ความถูกต้อง (Accuracy) 97.36%, ความแม่นยำ (Precision) 97.42%, ความไว (Recall) 97.36%, และค่า F1-score 97.34% ซึ่งสูงกว่าโมเดลซึ่งสูงกว่าโมเดลการวิเคราะห์อารมณ์ป้ายเดียวที่ดีที่สุดในงานวิจัยก่อนหน้านี้ซึ่งมีค่าทั้งสี่ตัวชี้วัดอยู่ที่ 89.96% (Khamphakdee & Seresangtakul, 2023)

จากผลการทดลองนี้แสดงให้เห็นว่า เทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) มีความสามารถในการช่วยให้โมเดลเข้าใจบริบทของข้อความได้ดีกว่าเทคนิคคลังคำศัพท์ (BoW) โดยการลดอิทธิพลของคำที่ปรากฏบ่อยแต่มีความสำคัญต่ำ และเพิ่มน้ำหนักให้กับคำที่มีความจำเพาะสูง ซึ่งสอดคล้องกับลักษณะของข้อความรีวิวที่มักมีคำหลากหลายและมีบริบทเฉพาะ นอกจากนี้ โมเดลที่ใช้การสุ่มป่าไม้ (Random Forest: RF) เมื่อเปรียบเทียบกับโมเดลอื่นๆ ที่ใช้ในการวิเคราะห์อารมณ์หลายป้าย เช่นการถดถอยโลจิสติกส์ (LR) และซัพพอร์ตเวกเตอร์แมชชีน (SVM) พบว่า โมเดลที่ใช้การสุ่มป่าไม้ (Random Forest: RF) มีความสามารถที่ดีกว่าในการวิเคราะห์อารมณ์หลายป้ายเนื่องจากสามารถเรียนรู้ลักษณะที่ไม่เป็นเชิงเส้นของข้อมูลได้ดี และมีความทนทานต่อปัญหาโอเวอร์ฟิตติง (Overfitting) มากกว่าโมเดลเชิงเส้นโดยทั่วไป ดังนั้น จากผลการทดลองทั้งหมดสามารถสรุปได้ว่า การประยุกต์ใช้โมเดลที่ใช้การสุ่มป่าไม้ (Random Forest: RF) ร่วมกับเทคนิคความถี่ของคำ-ส่วนกลับความถี่ของเอกสาร (TF-IDF) ถือเป็นแนวทางที่มีประสิทธิภาพสูงสุดในการวิเคราะห์อารมณ์ของรีวิวร้านอาหารภาษาไทย และมีศักยภาพในการประยุกต์ใช้ในงานด้านการทำความเข้าใจความคิดเห็น (Opinion Mining) และการประมวลผลภาษาธรรมชาติ (NLP) สำหรับภาษาไทย ได้อย่างมีประสิทธิภาพในอนาคต

เอกสารอ้างอิง

- Kemp, S. (2024). *Digital 2024 Global Overview Report*. Data Reportal. Retrieved from <https://datareportal.com/reports/digital-2024-global-overview-report>.
- Khamphakdee, N., & Seresangtakul, P. (2021). Sentiment Analysis for Thai Language in Hotel Domain Using Machine Learning Algorithms. *Acta Informatica Pragensia 2021*, 10(2), 155-171.
- Khamphakdee, N., & Seresangtakul, P. (2023). An Efficient Deep Learning for Thai Sentiment Analysis. *Data*, 8(5), 90.
- Liu, B. (2022). *Sentiment analysis and opinion mining*. Springer Nature.
- Liu, S., & Chen, J. (2014). A multi-label classification based approach for sentiment classification. *Expert Systems with Applications*, 42(3), 1083-1093.
- Lovert, O. (2024). *The impact of customer reviews on online sales*. Polaris Growth. Retrieved from <https://www.polarisgrowth.com/en/blog/the-impact-of-customer-review-on-online-sales>.
- Lowphansirikul, L., Polpanumas, C., Jantrakulchai, N., & Nutanong, S. (2021). Wangchanberta: Pretraining transformer-based thai language models. *arXiv preprint arXiv:2101.09635*.
- Marukatat, R. (2020). A Comparative Study of Using Bag-of-Words and Word-Embedding Attributes in the Spoiler Classification of English and Thai Text. *Applied Computing and Information Technology*, 847(1860-9503), 81-93.
- Pang, B., Lee, L., & Vaithyanathan, S. (2002). Thumbs up? Sentiment Classification using Machine Learning Techniques. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP 2002)*, pages 79-86. Association for Computational Linguistics.
- Phatthiyaphaibun, W., Chaovavanich, K., Polpanumas, C., Suriyawongkul, A., Lowphansirikul, L., & Chormai, P. (2019). *PyThaiNLP 2.1-dev7*. Retrieved from <https://zenodo.org/records/3519355>.
- Phatthiyaphaibun, W., Chaovavanich, K., Polpanumas, C., Suriyawongkul, A., Lowphansirikul, L., Chormai, P., Limkonchotiwat, P., Suntornitip, T., Udomcharoenchaikit, C. (2023). PyThaiNLP: Thai natural language processing in Python. *arXiv preprint arXiv:2312.04649*.
- Phuttakij, P., Plubin, B., Bunyatisai, W., Mouktonglang, T., & Plubin, S. (2025). BERT for Sentiment Analysis in Thai Hotel Reviews. *International Journal of Innovative Science and Research Technology (IJISRT)*, 10(2), 1774-1782.
- Puttipornchai, C., Chanyachatchawan, S., & Tuaycharoen, N. (2022). Multi-Label Classification for Articles in Thai Journal Database from Article's Abstract. *International Joint Conference on Computer Science and Software Engineering (JCSSE)*, 1-6.
- Sangsavate, S., Sinthupinyo, S. & Chandrachai, A. (2023). Experiments of Supervised Learning and Semi-Supervised Learning in Thai Financial News Sentiment: A Comparative Study. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 22(7), 1-36.
- Suciati, A., & Budi, I. (2019). Aspect-based Opinion Mining for Code-Mixed Restaurant Reviews in Indonesia. *International Conference on Asian Language Processing (IALP)*, 59-64.
- Tao, J., Fang, X. (2020). Toward multi-label sentiment analysis: a transfer learning based approach. *Journal of Big Data*, 7(1).

Thiengburanathum, P., & Charoenkwan, P. (2023). SETAR: Stacking Ensemble Learning for Thai Sentiment Analysis Using RoBERTa and Hybrid Feature Representation. *IEEE Access*, 11, 92822-92837.

Data Availability Statement: The raw data supporting the conclusions of this article will be made available by the authors, without undue reservation.

Conflicts of Interest: The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.



Copyright: © 2025 by the authors. This is a fully open-access article distributed under the terms of the Attribution-NonCommercial-NoDerivatives 4.0 International (CC BY-NC-ND 4.0).