

การวิเคราะห์ความคิดเห็นจากข้อความความคิดเห็นด้วยการเรียนรู้เชิงลึก

คณินาถ พงศ์ศักดิ์ศรี

TNII

วิทยานิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตรปริญญาวิศวกรรมศาสตรมหาบัณฑิต

บัณฑิตศึกษา สาขาเทคโนโลยีวิศวกรรม

สถาบันเทคโนโลยีไทย-ญี่ปุ่น

ปีการศึกษา 2566

SENTIMENT ANALYSIS FROM COMMENTS BY DEEP LEARNING

Kaninard Pongsaksri

TNII

A Thesis Submitted in Partial Fulfillment of Requirements
for The Degree of Master of Engineering Program in Engineering Technology
Graduate Studies

Thai-Nichi Institute of Technology
Academic Year 2023

หัวข้อวิทยานิพนธ์

การวิเคราะห์ความคิดเห็นจากข้อความความคิดเห็นด้วยการ
เรียนรู้เชิงลึก

โดย

คณินาถ พงศ์ศักดิ์ศรี

สาขาวิชา

เทคโนโลยีวิศวกรรม

อาจารย์ที่ปรึกษาวิทยานิพนธ์

ดร.อรรดา เชนโต๊ะ

บัณฑิตศึกษา สถาบันเทคโนโลยีไทย-ญี่ปุ่น อนุมัติให้บัณฑิตวิทยานิพนธ์ฉบับนี้เป็นส่วนหนึ่ง
ของการศึกษาตามหลักสูตรปริญญาวิทยาศาสตรบัณฑิต

.....รองอธิการบดีฝ่ายวิชาการ

(รองศาสตราจารย์ ดร.วรากร ศรีเชวงทรัพย์)

วันที่.....เดือน.....พ.ศ.....

คณะกรรมการสอบวิทยานิพนธ์

.....ประธานกรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.วิมล แสนอ้อม)

.....กรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.กัณติชา กิตติพิรัช)

.....กรรมการ

(ดร.ชัยชัย วรรณบุญ)

.....อาจารย์ที่ปรึกษาวิทยานิพนธ์

(ดร.อรรดา เชนโต๊ะ)

คณินาถ พงศ์ศักดิ์ศรี : การวิเคราะห์ความคิดเห็นจากข้อความความคิดเห็นด้วยการเรียนรู้เชิงลึก. อาจารย์ที่ปรึกษา : ดร. อัดนา เซนโตะ, 52 หน้า.

ในปัจจุบันการวิเคราะห์ข้อมูลจากความคิดเห็นของประชาชนบนโซเชียลเน็ตเวิร์ก เช่น Twitter หรือ Facebook หรืออื่นๆ ได้กลายมามีบทบาทสำคัญสำหรับการนำไปใช้งานด้านต่างๆ เช่น การตลาด การเมือง ฯลฯ โดยเฉพาะอย่างยิ่งงานด้านการพยากรณ์อารมณ์ความรู้สึกเชิงบวกและเชิงลบจากข้อความความคิดเห็นของผู้ที่ใช้งานโซเชียลเน็ตเวิร์กเพื่อใช้ในการปรับปรุงผลิตภัณฑ์หรือสินค้าต่างๆ นั้น เป็นที่ต้องการอย่างมาก ดังนั้นงานวิจัยนี้จึงนำเสนอโมเดลการเรียนรู้เชิงลึกแบบผสมผสานกับการประมวลผลภาษาธรรมชาติ (NLP) สำหรับการคัดแยกความรู้สึกเชิงบวกและเชิงลบ โดยในงานวิจัยนี้ได้ออกแบบขั้นตอนวิธี (อัลกอริทึม) ของการพยากรณ์ความรู้สึกจากข้อความวิจารณ์ที่ประกอบด้วยส่วนสำคัญอยู่ 2 ส่วนดังนี้ ส่วนแรกเป็นขั้นตอนของกระบวนการเตรียมข้อมูลซึ่งประกอบด้วยการทำความสะอาดข้อมูล และการดึงคุณลักษณะเด่นของข้อมูล หลังจากผ่านขั้นตอนการเตรียมข้อมูลแล้ว ข้อมูลข้อความจะอยู่ในรูปแบบของชุดข้อมูลที่มีค่าหรือประโยคแต่ละรายการแปลงเป็นเวกเตอร์ตัวเลข ในส่วนที่สองคือการออกแบบโมเดลสำหรับการเรียนรู้ด้วยการใช้โมเดลการเรียนรู้เชิงลึกแบบผสมผสานระหว่างอัลกอริทึม CNN และ LSTM เป็นการทำงานแบบคู่ขนานกันเพื่อนำจุดเด่นของทั้งสองมาใช้งานให้ได้ประโยชน์สูงสุด และในตอนท้ายของงานวิจัยนี้ก็ได้ออกแบบการทดลองเพื่อพิสูจน์ประสิทธิภาพของโมเดลอัลกอริทึมที่ได้นำเสนอด้วยค่าความแม่นยำที่ได้หลังทำการฝึกสอน (Accuracy) ซึ่งผลลัพธ์ของการทดลองโมเดลการเรียนรู้เชิงลึกที่นำเสนอก็สามารถคัดแยกความรู้สึกจากข้อความเฉลี่ยร้อยละ 95.84 จากตัวอย่างของการทดสอบทั้งหมด

บัณฑิตศึกษา

สาขาวิชา เทคโนโลยีวิศวกรรม

ปีการศึกษา 2566

ลายมือชื่อนักศึกษา

ลายมือชื่ออาจารย์ที่ปรึกษา

KANINARD PONGSAKSRI : SENTIMENT ANALYSIS FROM COMMENTS BY DEEP LEARNING. ADVISOR : DR.ADNA SENTO, 52 PP.

Analyzing data from public opinion on social networks such as Twitter, Facebook, etc., has become essential for applications in fields such as marketing, politics, etc. in particular, the task of forecasting the positive and the negative emotions from the user opinions on social networks to improve products or products has become highly in demand. Therefore, this research presents a deep learning model combined with natural language processing (NLP) for the classification of the positive and negative sentiments. In this research, an algorithm (algorithm) for sentiment classification from commentary text was designed that consists of two main parts as follows: The first part involves the data preparation process, including data cleaning and feature extraction. After the data preparation step, the textual data will be in the form of a dataset with each word or sentence transformed into numerical vectors. The second part is the design of the learning model using a mixed deep learning model combining CNN and LSTM algorithms, working in parallel to utilize the strengths of both for maximum benefit. At the end of this research, experimental designs were also conducted to prove the efficiency of the proposed algorithm models by evaluating the accuracy after training. The results of the deep learning model experiments presented can separate sentiment from text with an average accuracy of 95.84% from the entire test sample.

Graduate Studies

Field of Engineering of Technology

Academic Year 2023

Student's Signature.....

Advisor's Signature.....

กิตติกรรมประกาศ

วิทยานิพนธ์ฉบับนี้สำเร็จลุล่วงได้ด้วยดีโดยได้รับคำปรึกษาจากอาจารย์ที่ปรึกษา ดร.อัฒนา เซนโตะ จากประธานกรรมการสอบ รวมถึงอาจารย์ทุกท่านที่มีได้เอ่ยนามมา ณ ที่นี้ด้วย ที่ได้คอยให้คำแนะนำและให้ความช่วยเหลือในทุกขั้นตอนของการทำวิจัยครั้งนี้เป็นอย่างดี เพื่อให้งานวิจัยฉบับนี้ มีความเรียบร้อยสมบูรณ์

และขอขอบพระคุณพ่อ แม่ ครอบครัว และเพื่อนๆ ทุกท่านที่คอยให้กำลังใจให้คำปรึกษา และสนับสนุนการทำงานวิจัยจนประสบผลสำเร็จลุล่วงไปได้ด้วยดี

คณินาถ พงศ์ศักดิ์ศรี

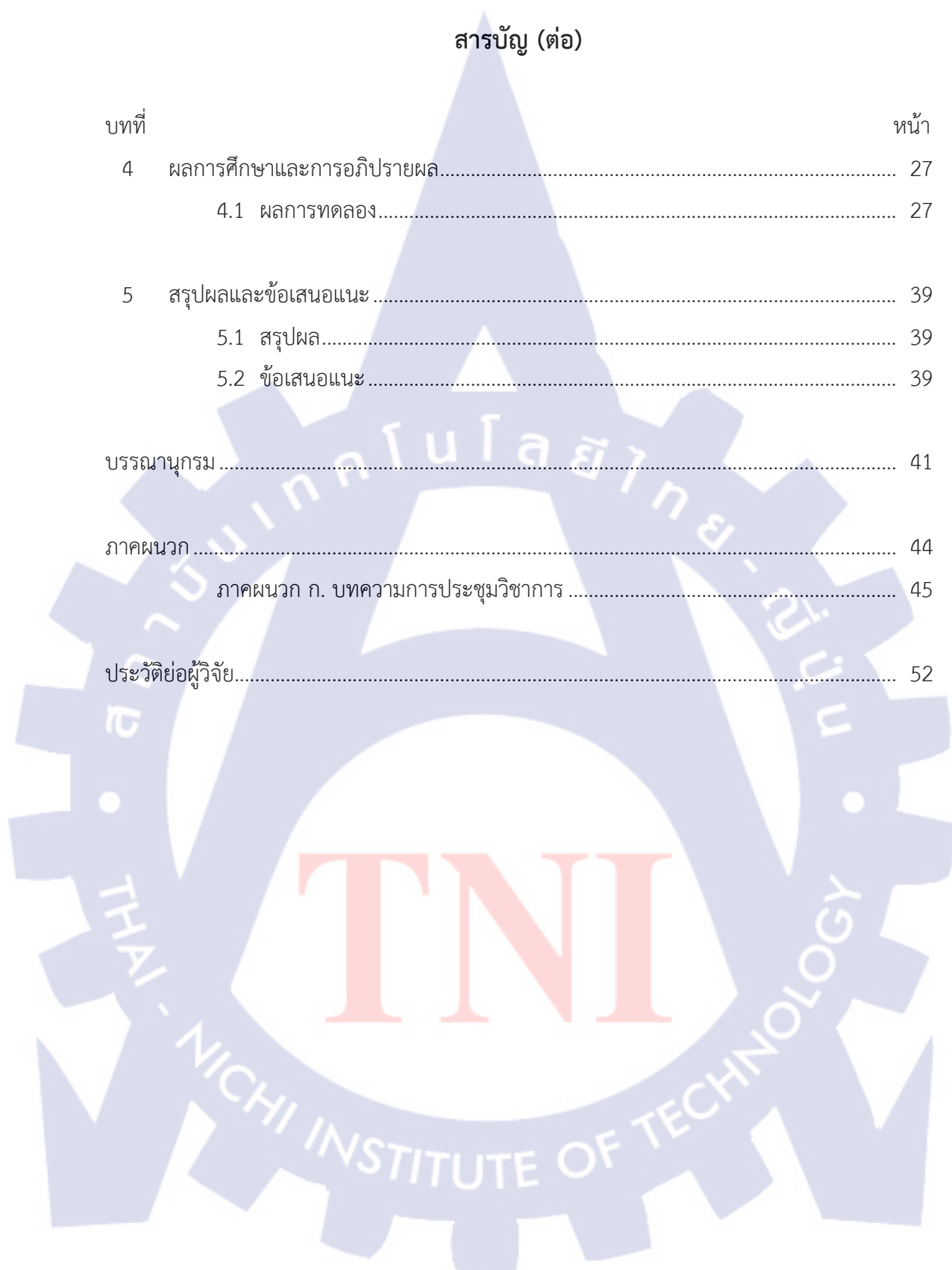


สารบัญ

	หน้า
บทคัดย่อภาษาไทย.....	ง
บทคัดย่อภาษาอังกฤษ.....	จ
กิตติกรรมประกาศ.....	ฉ
สารบัญ	ช
สารบัญตาราง	ฌ
สารบัญรูป.....	ญ
บทที่	
1 บทนำ	1
1.1 ที่มาและความสำคัญ.....	1
1.2 วัตถุประสงค์ของการวิจัย.....	2
1.3 ขอบเขตของงานวิจัย	2
1.4 ผลที่คาดว่าจะได้รับ	3
1.5 ขั้นตอนการวิจัย	3
1.6 แผนงานและระยะเวลาดำเนินการ	3
2 ทฤษฎีและงานวิจัยที่เกี่ยวข้อง.....	5
2.1 การเตรียมข้อมูล.....	6
2.2 งานวิจัยที่เกี่ยวข้อง	14
3 วิธีการวิจัย	19
3.1 กรอบความคิด	19
3.2 ขั้นตอนการทดลอง.....	24
3.3 การดำเนินการวิจัยและเก็บข้อมูล.....	25
3.4 ทำการทดลอง.....	26

สารบัญ (ต่อ)

บทที่	หน้า
4 ผลการศึกษาและการอภิปรายผล.....	27
4.1 ผลการทดลอง.....	27
5 สรุปผลและข้อเสนอแนะ.....	39
5.1 สรุปผล.....	39
5.2 ข้อเสนอแนะ.....	39
บรรณานุกรม.....	41
ภาคผนวก.....	44
ภาคผนวก ก. บทความการประชุมวิชาการ.....	45
ประวัติย่อผู้วิจัย.....	52



สารบัญตาราง

ตาราง	หน้า
1.1 แผนงานและระยะเวลาดำเนินการวิจัย	3
4.1 ร้อยละความถูกต้องของผลการทดลองการวิเคราะห์อารมณ์ความรู้สึกบวกและลบ	27
4.2 ค่า F1 Score ของผลการทดลองการวิเคราะห์อารมณ์ความรู้สึกบวกและลบ.....	38



สารบัญรูป

รูป	หน้า
2.1 แผนภาพแสดงการดำเนินงานของระบบการ.....	5
2.2 ตัวอย่างข้อมูลดิบ Sentiment140.....	6
2.3 ตัวอย่างข้อมูลดิบ Tweets Airline.....	7
2.4 ตัวอย่างข้อมูลดิบ บทวิจารณ์ภาพยนตร์ IMDB.....	7
2.5 Natural Language Process Pipeline	8
2.6 Word Embedding.....	9
2.7 Long Short-Term Memory	11
2.8 GRU	11
2.9 Convolutional Neural Network.....	12
2.10 Model CNN-Static	15
2.11 ผลการทดลองแต่ละแบบจำลอง.....	15
2.12 Model Architecture ของแบบจำลอง.....	15
2.13 Feed-Forward ของ CNN	16
2.14 Model Architecture ของแบบจำลอง.....	17
2.15 Layers ของแบบจำลอง M-Hybrid.....	17
2.16 ผลทดลองความแม่นยำของแบบจำลองพื้นฐาน และ M-Hybrid.....	18
3.1 RNN Model.....	19
3.2 LSTM Model.....	20
3.3 GRU Model.....	21
3.4 CNN-LSTM Model.....	22
3.5 Proposed Model.....	23
3.6 โครงสร้างของโมเดลการเรียนรู้เชิงลึกแบบผสมผสานที่นำเสนอ.....	24
3.7 ตัวอย่างข้อมูลดิบสำหรับการทดลอง.....	25
3.8 ขั้นตอนการดำเนินการวิจัย.....	25
3.9 ตัวอย่างข้อมูลดิบสำหรับการทดลอง.....	26
4.1 กราฟของการทดลองของอัลกอริทึม RNN ด้วยชุดข้อมูล Sentiment140	28
4.2 กราฟของการทดลองของอัลกอริทึม RNN ด้วยชุดข้อมูล Tweets Airline.....	28

สารบัญรูป (ต่อ)

รูป	หน้า
4.3 กราฟของการทดลองของอัลกอริทึม RNN ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	28
4.4 ตาราง Confusion Matrix ของอัลกอริทึม RNN ด้วยชุดข้อมูล Sentiment140.....	29
4.5 ตาราง Confusion Matrix ของอัลกอริทึม RNN ด้วยชุดข้อมูล Tweets Airline.....	29
4.6 ตาราง Confusion Matrix ของอัลกอริทึม RNN ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	29
4.7 กราฟของการทดลองของอัลกอริทึม LSTM ด้วยชุดข้อมูล Sentiment140.....	30
4.8 กราฟของการทดลองของอัลกอริทึม LSTM ด้วยชุดข้อมูล Tweets Airline.....	30
4.9 กราฟของการทดลองของอัลกอริทึม LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	30
4.10 ตาราง Confusion Matrix ของอัลกอริทึม LSTM ด้วยชุดข้อมูล Sentiment140.....	31
4.11 ตาราง Confusion Matrix ของอัลกอริทึม LSTM ด้วยชุดข้อมูล Tweets Airline	31
4.12 ตาราง Confusion Matrix ของอัลกอริทึม LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB	31
4.13 กราฟของการทดลองของอัลกอริทึม GRU ด้วยชุดข้อมูล Sentiment140	32
4.14 กราฟของการทดลองของอัลกอริทึม GRU ด้วยชุดข้อมูล Tweets Airline	32
4.15 กราฟของการทดลองของอัลกอริทึม GRU ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	32
4.16 ตาราง Confusion Matrix ของอัลกอริทึม GRU ด้วยชุดข้อมูล Sentiment140.....	33
4.17 ตาราง Confusion Matrix ของอัลกอริทึม GRU ด้วยชุดข้อมูล Tweets Airline.....	33
4.18 ตาราง Confusion Matrix ของอัลกอริทึม GRU ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	33
4.19 กราฟของการทดลองของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Sentiment140.....	34
4.20 กราฟของการทดลองของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Tweets Airline	34
4.21 กราฟของการทดลองของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	34
4.22 ตาราง Confusion Matrix ของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Sentiment140 ...	35
4.23 ตาราง Confusion Matrix ของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Tweets Airline...	35
4.24 ตาราง Confusion Matrix ของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	35

สารบัญรูป (ต่อ)

รูป	หน้า
4.25 กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วยชุดข้อมูล Sentiment140.....	36
4.26 กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วยชุดข้อมูล Tweets Airline.....	36
4.27 กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	36
4.28 ตาราง Confusion Matrix นำเสนอด้วยชุดข้อมูล Sentiment140.....	37
4.29 ตาราง Confusion Matrix นำเสนอด้วยชุดข้อมูล Tweets Airline	37
4.30 ตาราง Confusion Matrix นำเสนอด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB.....	37



บทที่ 1

บทนำ

1.1 ที่มาและความสำคัญ

จาก [1] ได้ให้ข้อมูลเกี่ยวกับพฤติกรรมของผู้บริโภคในยุคปัจจุบันที่ใช้ช่องทางการสั่งซื้อออนไลน์มากขึ้น จุดนี้เองทำให้เกิดการวิจารณ์สินค้าในกลุ่มผู้บริโภคนั้นๆ ที่เขียนลงบนโซเชียลมีเดียมากขึ้นตามไปด้วย ด้วยเหตุนี้ทำให้เกิดการตัดสินใจของผู้บริโภคต่อสินค้าใดๆ ก่อนการเลือกซื้อด้วยการอ่านคำวิจารณ์ของสินค้าเสมอ ในทางกลับกันผู้ผลิตสินค้าก็สามารถพัฒนาสินค้าของตนด้วยการนำคำวิจารณ์เหล่านั้นมาปรับปรุงและพัฒนางานสินค้าของตัวเองอีกด้วย และปัจจุบันก็เป็นที่ยอมรับกันอย่างกว้างขวางว่าการวิเคราะห์ความรู้สึกนั้นมีประโยชน์อย่างมากในการดำเนินกิจการด้านต่างๆ เช่น ระบบธุรกิจอัจฉริยะและอีคอมเมิร์ซ ซึ่งทำให้สามารถศึกษาความคิดเห็นของลูกค้าเพื่อให้การสนับสนุนลูกค้าที่ดีขึ้น สร้างผลิตภัณฑ์ที่ดีขึ้น หรือปรับปรุงกลยุทธ์ทางการตลาดเพื่อดึงดูดลูกค้าใหม่ สามารถใช้การวิเคราะห์ความรู้สึกเพื่อสรุปความคิดเห็นของผู้ใช้เกี่ยวกับเหตุการณ์หรือผลิตภัณฑ์ ผลลัพธ์ของการวิเคราะห์สามารถช่วยให้ได้รับข้อมูลเชิงลึกมากขึ้นเกี่ยวกับความสนใจหรือความคิดเห็นของลูกค้าเกี่ยวกับแนวโน้มอุตสาหกรรม ส่งผลให้เกิดงานวิจัยที่มุ่งเน้นเกี่ยวกับงานการสร้างแบบจำลองคุณภาพสูงเพื่อแก้ปัญหาความซับซ้อนของข้อมูลขนาดใหญ่ที่มีแนวโน้มที่เพิ่มขึ้นอย่างต่อเนื่อง เพื่อการพัฒนางานด้านการนำผลจากการศึกษาไปใช้ในส่วนของการขยายการวิเคราะห์ความรู้สึกไปสู่การใช้งานที่หลากหลายไม่ว่าจะเป็นด้านสุขภาพ [2, 3] ด้านการเงิน [4] ด้านการวิเคราะห์การตลาด [5, 6] และอื่นๆ [7, 8, 9, 10] ยิ่งไปกว่านั้นเมื่อไม่นานมานี้ การนำแนวคิดนี้สร้างงานวิจัยเกี่ยวกับการวิเคราะห์พฤติกรรมของผู้บริโภคจากความคิดเห็นหรือความวิจารณ์แบบอัตโนมัติมากขึ้นเรื่อยๆ เช่นกัน ตัวอย่างเช่น [11]

โดยมีโมเดลของการเรียนรู้เชิงลึกเป็นตัวขับเคลื่อนให้เกิงานวิจัยที่มีประสิทธิภาพมากขึ้นซึ่งส่วนใหญ่ของงานวิจัยได้นำ Convolutional Neural Networks (CNN) หรือ Recurrent Neural Networks (RNN) หรือ Long Short-Term Memory (LSTM) ที่ให้ผลลัพธ์ความแม่นยำโดยรวมที่ค่อนข้างสูงเพื่อให้งานวิจัยได้รับการยอมรับมากขึ้น ด้วยเหตุผลนี้เอง ในงานวิจัยนี้จึงตระหนักเห็นถึงความจำเป็นอย่างยิ่งที่จะต้องพัฒนางานวิจัยในด้านนี้เพื่อเพิ่มและปรับปรุงประสิทธิภาพในการเรียนรู้ความคิดเห็นของประชากรผู้บริโภคในลักษณะการสั่งซื้อออนไลน์ หรือบนพื้นที่โซเชียล โดยกระบวนการวิธีของงานวิจัยนี้ได้้นำการประมวลผลภาษาธรรมชาติ (NLP) มาผนวกกับอัลกอริทึมของโมเดลการเรียนรู้เชิงลึกเพื่อแก้ปัญหการวิเคราะห์ความรู้สึกสำหรับการคัดแยกความรู้สึกเชิงบวก และเชิงลบซึ่งประกอบด้วยขั้นตอนของกระบวนการเตรียมข้อมูลซึ่งประกอบด้วยการทำความสะอาด โดยหลังจากเตรียมชุดข้อมูลเป็นที่เรียบร้อยแล้ว ประการต่อมาคือการแปลงข้อความวิจารณ์เป็นเวกเตอร์ (หรือที่เรียกว่า Feature Vector)

โดยใช้วิธี Word Embedding ซึ่งทำให้ได้ผลลัพธ์เวกเตอร์คุณลักษณะที่พร้อมจะนำไปป้อนให้กับอัลกอริทึมการเรียนรู้เชิงลึกต่อไป โดยในส่วนของท้ายของการวิจัยนี้ได้ออกแบบการทดลองเพื่อให้ผู้อ่านมั่นใจในระบบที่ได้นำเสนอขึ้นมากขึ้นด้วยการนำเมตริกความน่าเชื่อถือ (Confusion Matrix) เพื่อประเมินความสามารถของโมเดล นอกจากนี้ยังการทดลองเพื่อประเมินประสิทธิภาพของโมเดลการเรียนรู้เชิงลึกและเทคนิคที่เกี่ยวข้องในชุดข้อมูลเกี่ยวกับหัวข้อต่างๆ เพื่อเปรียบเทียบวิธีการวิเคราะห์ความรู้สึกกับงานวิจัยอื่นๆ ซึ่งจากการทดลองนี้ได้ชี้ให้เห็นว่าสามารถนำคำวิจารณ์ที่มีอยู่ในปัจจุบันมาวิเคราะห์และคัดแยกความรู้สึกเชิงบวกและลบจากข้อความเฉลี่ยในระดับที่ดีมากจากตัวอย่างของการทดสอบทั้งหมด ผู้อ่านสามารถเปิดอ่านงานวิจัยนี้ในหัวข้ออื่นๆ ได้ในลำดับถัดจากนี้คือ ส่วนที่สองจะกล่าวถึงระบบที่นำเสนอจากนั้นจะเป็นหัวข้อของการทดลอง และสุดท้ายจะเป็นหัวข้อสรุปและวิจารณ์ผลการทดลอง

1.2 วัตถุประสงค์ของการวิจัย

1.2.1 สามารถพัฒนาแบบจำลองที่สามารถแบ่งประเภทอารมณ์ของข้อความได้ด้วยความแม่นยำ (Accuracy) ที่สูงกว่า 79%

1.2.2 สามารถใช้โปรแกรมภาษา Python ในการสร้างแบบจำลองการเรียนรู้เชิงลึก (Deep Learning) หรือโครงข่ายประสาทเทียม (Neuron Network) สำหรับแยกแยะอารมณ์ของข้อความ เป็นเชิงบวกและเชิงลบ ด้วยแบบจำลอง LSTM

1.2.3 เพื่อพัฒนาแนวทางการแยกแยะรูปแบบข้อความเชิงบวก (Positive) และเชิงลบ (Negative) ด้วยแบบจำลอง Deep Learning

1.3 ขอบเขตของงานวิจัย

1.3.1 แบบจำลองนี้ถูกพัฒนาขึ้นโดยเครื่องคอมพิวเตอร์ประสิทธิภาพสูงที่ได้ติดตั้งระบบปฏิบัติการ Windows 11 Professionnal ชนิด 64 บิต มีหน่วยความจำ 32 GB และมีพื้นที่เก็บข้อมูลขนาด 1 TB มีการใช้ตัวประมวลผลกราฟิกรุ่น RTX 3090 จากบริษัท Nvidia ซึ่งมีหน่วยความจำ 24 GB

1.3.2 แบบจำลองในงานวิจัยนี้สามารถแยกแยะข้อความต่างๆ ได้ 2 แบบ คือ เชิงบวก และเชิงลบ

1.3.3 แบบจำลองในงานวิจัยถูกทดสอบเพื่อหาความแม่นยำ (Validation Accuracy) ด้วยข้อมูลอีกชุดหนึ่งที่ได้จาก Ssocial Media

1.3.4 แบบจำลองในงานวิจัยนี้สามารถใช้ได้กับข้อมูลที่สมบูรณ์ หรือไม่มีการพิมพ์ที่ตกหล่นตัวอักษรบางตัวไป

ตารางที่ 1.1 แผนงานและระยะเวลาดำเนินการวิจัย (ต่อ)

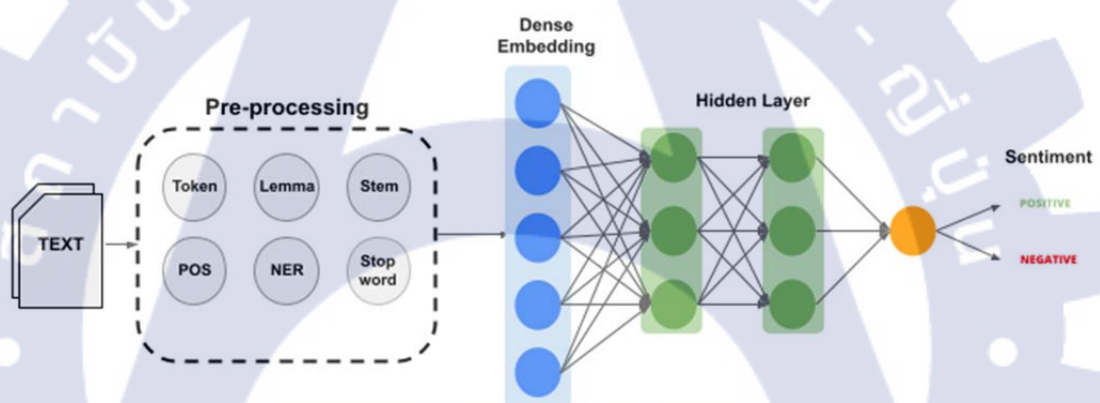
การดำเนินการ	2022				2023												2024		
	9	10	11	12	1	2	3	4	5	6	7	8	9	10	11	12	1	2	3
5. ปรับแต่งและพัฒนา โครงสร้างแบบจำลอง																			
6. ทดสอบแบบจำลอง ครั้งที่ 2																			
7. ทดสอบแบบจำลอง ครั้งที่ 3																			
8. วิเคราะห์และสรุปผล การทดลอง																			



บทที่ 2

ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

ปัจจุบันอัลกอริทึมสำหรับการรู้จำ คัดแยก และการทำนายที่ได้รับความนิยมอย่างกว้างขวาง คือการเรียนรู้เชิงลึกที่มีความสามารถเหนือกว่าอัลกอริทึมดั้งเดิม (การเรียนรู้ของเครื่อง เช่น Support Vector Machine (SVM), เครือข่ายแบบเบย์ หรือแผนผังการตัดสินใจ) เนื่องจากได้รับความนิยมอย่างมาก มันทำให้เกิดการพัฒนาสิ่งใหม่ๆ ในโมเดลการเรียนรู้เชิงลึก คุณลักษณะต่างๆ จะได้รับการเรียนรู้และดึงข้อมูลโดยอัตโนมัติ ทำให้ได้ความแม่นยำและประสิทธิภาพที่ดีขึ้น โดยทั่วไปการเรียนรู้เชิงลึก เครือข่ายประสาทเทียมและการเรียนรู้เชิงลึกในปัจจุบันให้แนวทางแก้ไขที่ดีที่สุดสำหรับปัญหามากมายในด้านการรู้จำภาพและเสียง เช่นเดียวกับการประมวลผลภาษาธรรมชาติ โดยในหัวข้อนี้กล่าวถึงเทคนิคการเรียนรู้เชิงลึกซึ่งมีรายละเอียดดังรูปที่ 2.1



รูปที่ 2.1 แผนภาพแสดงการดำเนินงานของระบบการ

ส่วนของประกอบของโมเดลการเรียนรู้เชิงลึกเพื่อการประยุกต์ใช้งานด้านการประมวลผลทางภาษา (NLP) ได้แก่ ส่วนของกระบวนการจัดเตรียมข้อมูล (Pre-Processing) กระบวนการจัดเรียงคำ (Dense Embedding) ส่วนของโครงข่ายประสาทเทียม (Neural Network) และส่วนของเอาต์พุตใน ซึ่งมีรายละเอียดดังต่อไปนี้

2.1 การเตรียมข้อมูล

2.1.1 ชุดข้อมูล (Dataset)

2.1.1.1 Sentiment140 จาก Stanford University [12]

เป็นข้อความทวิตเกี่ยวกับผลิตภัณฑ์จำนวน 1.5 ล้านกว่าทวิตและได้ระบุค่าความรู้สึกของผู้วิจารณ์ด้านลบ และด้านบวก

	review	sentiment
1	This is June 17 but where's the Iphone 3.0	0
2	Hope to see the Peace in my Homeland IR...	0
3	sleeping alone tonight	0
4	Chemistry again	0
5	for future reference: never EVER read plot...	0
6	I really don't like the fact that Leah will be ...	0
7	I want an iPhone reallyyy badly	0
8	can no longer remember what it was like t...	0
9	is going crazy	0
10	is sick as in in bed and not at work	0

รูปที่ 2.2 ตัวอย่างข้อมูลดิบ Sentiment140

2.1.1.2 Tweets Airline [13]

เป็นชุดข้อมูลทวิตที่มีความคิดเห็นของผู้ใช้เกี่ยวกับสายการบินของสหรัฐฯ มีการรวบรวมข้อมูลในเดือนกุมภาพันธ์ 2015 มีตัวอย่าง 1 หมื่นกว่าตัวอย่าง และแบ่งออกเป็นความรู้สึกเชิงลบ เป็นกลาง และเชิงบวก ซึ่งในการทดสอบนี้ได้ตัดแบ่งนำมาใช้เฉพาะส่วนของข้อมูลการวิจารณ์ที่มีความรู้สึกเป็นบวกและลบเท่านั้น

	review	sentiment
1	@VirginAmerica it's really aggressive to bla...	negative
2	@VirginAmerica and it's a really big bad thi...	negative
3	@VirginAmerica seriously would pay \$30 a ...	negative
4	@VirginAmerica SFO-PDX schedule is still ...	negative
5	@VirginAmerica I flew from NYC to SFO la...	negative
6	@VirginAmerica why are your first fares in ...	negative
7	@VirginAmerica you guys messed up my s...	negative
8	@VirginAmerica status match program. I a...	negative
9	@VirginAmerica What happened 2 ur vega...	negative
10	@VirginAmerica amazing to me that we can...	negative

รูปที่ 2.3 ตัวอย่างข้อมูลดิบ Tweets Airline

2.1.1.3 บทวิจารณ์ภาพยนตร์ IMDB ได้มาจาก Stanford University [14]

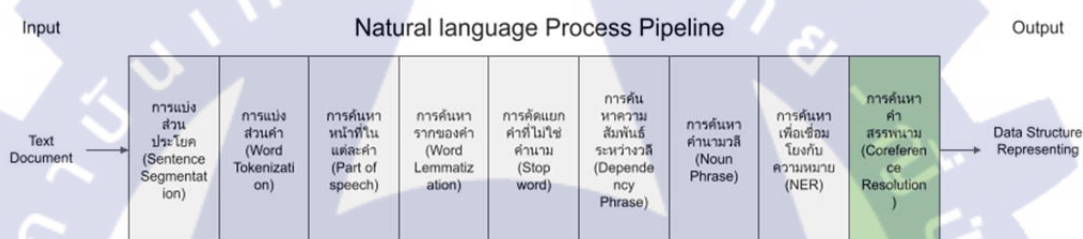
ชุดข้อมูลนี้มีความคิดเห็นจากผู้ชมเกี่ยวกับเรื่องราวของภาพยนตร์ มี 50,000 กว่าตัวอย่าง ซึ่งแบ่งเป็นบวกและลบเช่นกัน

	review	sentiment
1	One of the other reviewers has mentioned that after watchi...	positive
2	A wonderful little production. The filming techni...	positive
3	I thought this was a wonderful way to spend time on a too ...	positive
4	Basically there's a family where a little boy (Jake) thinks th...	negative
5	Petter Matte's "Love in the Time of Money" is a visually stu...	positive
6	Probably my all-time favorite movie, a story of selflessness,...	positive
7	I sure would like to see a resurrection of a up dated Seahu...	positive
8	This show was an amazing, fresh & innovative idea in the 7...	negative
9	Encouraged by the positive comments about this film on he...	negative
10	If you like original gut wrenching laughter you will like this ...	positive

รูปที่ 2.4 ตัวอย่างข้อมูลดิบ บทวิจารณ์ภาพยนตร์ IMDB

2.1.2 กระบวนการจัดเตรียมข้อมูล (Pre-Processing)

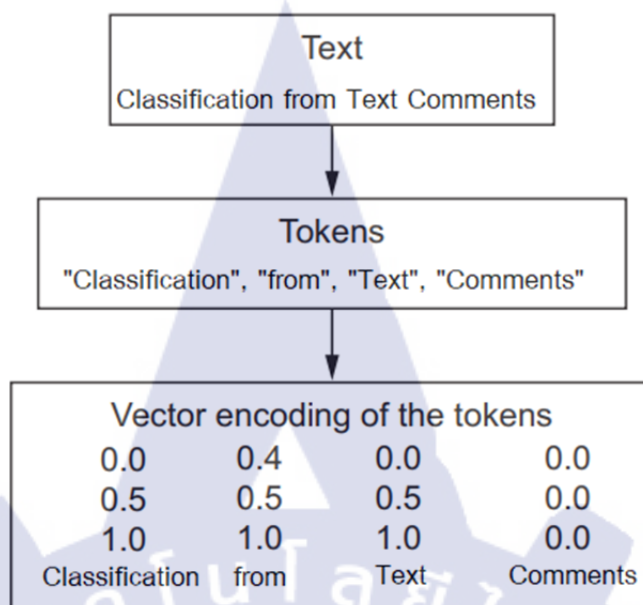
หัวข้อนี้จะกล่าวถึงขั้นตอนของการจัดเตรียมข้อมูล (Dataset) ก่อนการนำไปป้อนให้กับโมเดลการเรียนรู้เชิงลึกเพื่อคัดแยกความรู้สึก ในอดีตที่ผ่านการเตรียมข้อมูลมักจะมีขั้นตอนอย่างง่าย ได้แก่ การโหลดข้อมูล ทำความสะอาดข้อมูล การทำระบุค่าเป็นตัวเลข (Bag-of-Word) แต่ในปัจจุบันได้มีการพัฒนาขั้นตอนนี้เพื่อให้ประสิทธิภาพของการนำไปใช้สูงขึ้นจึงได้ปรับเปลี่ยนขั้นตอนสำหรับการเตรียมข้อมูลดังต่อไปนี้ การแบ่งส่วนประโยค (Sentence Segmentation) → การแบ่งส่วนคำ (Word Tokenization) → การค้นหาหน้าที่ในแต่ละคำ (Part of Speech) → การค้นหารากของคำ (Word Lemmatization) → การคัดแยกคำที่ไม่ใช่คำนาม (Stop Word) → การค้นหาความสัมพันธ์ระหว่างวลี (Dependency Phrase) → การค้นหาคำนามวลี → การค้นหาเพื่อเชื่อมโยงกับความหมาย (NER) → การค้นหาคำสรพนาม ซึ่งสามารถแสดงดังรูปที่ 2.5



รูปที่ 2.5 Natural Language Process Pipeline

2.1.3 กระบวนการจัดเรียงคำ (Dense Embedding)

เมื่อได้เตรียมคำ (Word) เป็นที่เรียบร้อยแล้ว จากนั้นให้กำหนดค่าต่างด้วยตัวเลขตามลำดับของคำในประโยคเพื่อสร้างเป็นอินพุตให้กับโครงข่ายประสาทเทียมมีรายละเอียดดังต่อไปนี้ การวิเคราะห์โมเดลการเรียนรู้ของเครื่องเรื่องของการประมวลผลทางภาษาจำเป็นต้องแปลงข้อความให้เป็นตัวเลขเสียก่อน โดยทั่วไปสามารถเปรียบเทียบกับตัวแปรหมวดหมู่ได้ด้วยการเข้ารหัสแบบ One-Hot เพื่อแปลงคุณลักษณะที่เป็นหมวดหมู่เป็นตัวเลข ซึ่งการเข้ารหัสแบบ One-Hot กับคำในข้อมูลที่เป็นข้อความนั้นทำให้ต้องสร้างคุณลักษณะที่ผ่านการเข้ารหัสเท่ากับจำนวนข้อความในชุดข้อมูลตัวอย่างเช่นหากมีคำศัพท์ 10,000 คำ จะได้ค่าคุณลักษณะ 10,000 รายการทำให้ต้องมีพื้นที่จัดเก็บขนาดใหญ่และประสิทธิภาพของโมเดลลดต่ำลง เพราะฉะนั้นจำเป็นต้องใช้ขั้นตอนวิธีสำหรับการกำหนดค่าคุณลักษณะแทนข้อความโดยในที่นี้ขอเสนอวิธีการที่มีชื่อว่า Word Embedding ซึ่งเป็นการแปลงแต่ละคำเป็นเวกเตอร์ทำให้ช่วยเพิ่มประสิทธิภาพและลดขนาดข้อมูลลงดังรูปที่ 2.6



รูปที่ 2.6 Word Embedding

โดยทั่วไป การทำ Word Embedding เป็นการแปลงแต่ละคำเป็นเวกเตอร์ตามขนาดที่กำหนดได้ เวกเตอร์ผลลัพธ์เป็นเวกเตอร์สัมพันธ์ กล่าวคือแทนค่าความสัมพันธ์ระหว่างคำโดยมีขั้นตอนเริ่มจากการกำหนดขนาดและแทนค่าตัวเลขของคำทั้งหมดโดยไม่ซ้ำกัน หากแต่ละประโยคของข้อมูลมีจำนวนคำในประโยคไม่เท่ากันให้เพิ่ม 0 (การทำ Padding) เพื่อให้ได้เวกเตอร์ขนาดเท่ากัน หลังจากนั้นให้กำหนดขนาดเวกเตอร์ผลลัพธ์ที่ต้องการ โดยทั่วไป Word Embedding สามารถสร้างขึ้นได้โดยใช้วิธีการต่างๆ เช่น โครงข่ายประสาทเทียม เมทริกซ์แบบพิเศษ แบบจำลองความน่าจะเป็น ฯลฯ Word2Vec สามารถทำได้ในรูปแบบ Continuous Bag of Word (CBOW) และ Skip-gram โดย CBOW เป็นโมเดลที่ทำนายคำปัจจุบันที่กำหนดคำบริบทภายในลำดับคำนั้นๆ โครงข่ายของ CBOW ประกอบด้วยอินพุตที่เป็นคำที่อยู่ลำดับใกล้ๆ กัน และเอาต์พุตกำหนดให้เป็นคำปัจจุบัน ส่วนชั้นซ่อน (Hidden Layer) เป็นตัวกำหนดมิติที่เราต้องการแสดงคำปัจจุบัน สำหรับ Skip-Gram จะเป็นโมเดลที่ทำนายคำโดยรอบ โดยโครงสร้างเลเยอร์ของ Skip-Gram ประกอบด้วยอินพุต (คำปัจจุบัน) และเลเยอร์เอาต์พุตกำหนดให้เป็นคำโดยรอบ ส่วนชั้นซ่อน (Hidden Layer) เป็นตัวกำหนดมิติที่เราต้องการแสดงคำปัจจุบัน

2.1.4 โครงข่ายประสาทเทียม (Neural Network) และส่วนของเอาต์พุต

ในส่วนของการคัดแยกกลุ่มข้อความจะใช้โครงข่ายประสาทเทียมแบบวนซ้ำหรือที่เรียกว่า Recurrent Neural Network (RNN) เนื่องจากเป็นโครงข่ายที่หน้าทีในการประมวลผลข้อมูลตามลำดับทำให้สามารถจดจำการคำนวณข้อมูลก่อนหน้านี้ ปัจจุบัน RNN ได้รับความนิยมจึงมีการ

พัฒนาต่อยอดทำให้เกิดโมเดลใหม่ๆ เช่น โครงข่าย Long-Short Term Memory (LSTM) หรือจะเป็น GRU เป็นต้นนอกจากนี้ยังมีงานวิจัย เกี่ยวกับการใช้งานทำนายข้อความจากความรู้สึกซึ่งมักจะใช้โครงข่าย LSTM เพื่อคัดแยกข้อความดังรูปที่ 2.6 ซึ่งประกอบด้วยชั้นอินพุต (ชุดข้อมูลที่เป็นข้อความ) ชั้นของ Embedding Matrix ซึ่งเป็นส่วนชั้นที่ต้องรับข้อมูลอินพุตผ่านการเตรียมข้อมูลแล้ว โดยชั้นเลเยอร์นี้ จะต้องกำหนดให้มีตัวกรองเพื่อคัดคุณลักษณะพิเศษตามหลักการของโครงข่ายประสาทเทียมแบบ Convolutional Neural Network (CNN) เมื่อผ่านกระบวนการนี้แล้วชั้นถัดไปจะเป็นชั้น LSTM และ ส่วนสุดท้ายจะเป็นโครงข่ายแบบเชื่อมต่อกันหมด (Fully Connected Network) ที่มีชั้นสุดท้ายเป็น คลาสที่ต้องทำนายกำหนดให้เป็นความรู้สึกด้านบวก และความรู้สึกลบ

2.1.4.1 Natural Language Processing

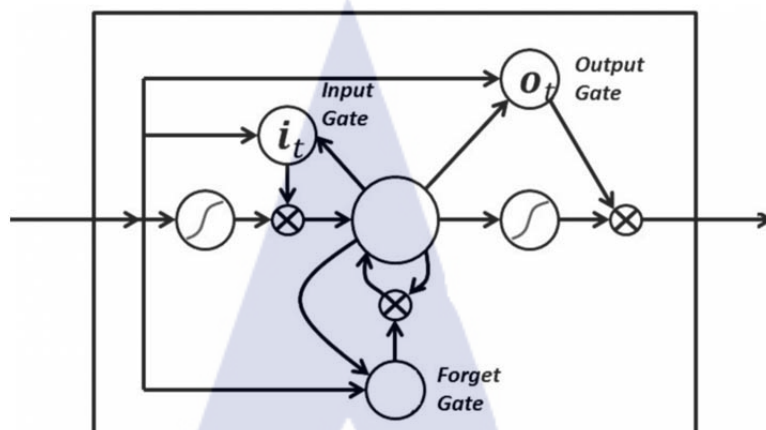
Natural Language Processing หรือ NLP นั้นถูกจำกัดความไว้อย่างกว้าง ว่าเป็นการประมวลผลภาษาธรรมชาติเช่นบทสนทนาหรือข้อความโดยอัตโนมัติด้วยโปรแกรม

ภาษาธรรมชาติ หมายถึง วิธีการที่มนุษย์ใช้สื่อสารกันเช่นบทสนทนาหรือ ข้อความซึ่งจำเป็นต้องมีกระบวนการในการตีความและใช้ตรรกะเช่นเดียวกับการจัดการข้อมูลแบบอื่น โดยการประมวลผลภาษาธรรมชาติผสมผสานภาษาศาสตร์เชิงคำนวณและการสร้างแบบจำลองตาม กฎเกณฑ์ของภาษามนุษย์เข้ากับแบบจำลองทางสถิติ Machine Learning และแบบจำลองการเรียนรู้ เชิงลึกซึ่งเทคโนโลยีเหล่านี้ร่วมกันทำให้คอมพิวเตอร์สามารถประมวลผลภาษามนุษย์ในรูปแบบของ ข้อมูลข้อความหรือเสียง และเพื่อเข้าใจความหมายที่สมบูรณ์ด้วยเจตนาและความรู้สึกของผู้พูดหรือ ผู้เขียน

2.1.4.2 Long Short-Term Memory

Long Short-Term Memory หรือ LSTM คือโครงข่ายประสาทเทียมแบบ วงกลับชนิดหนึ่งที่ถูกพัฒนาต่อยอดมาจาก RNN หรือ Recurrent Neural Network ซึ่งหลักการทำงานของ RNN คือนำผลลัพธ์ที่ได้จากการคำนวณย้อนกลับมาใช้เป็นข้อมูลเข้าอีกครั้งทำให้ลักษณะ เด่นของ RNN คือสามารถใช้ข้อมูลก่อนหน้าในอดีตทำนายสิ่งที่อาจเกิดขึ้นในอนาคตได้ หากแต่ RNN มี ข้อเสียคือสามารถใช้ข้อมูลย้อนหลังได้เพียงช่วงระยะเวลาสั้นๆ สืบเนื่องมาจากค่าจากการคำนวณย้อนหลัง ที่จะลดลงเมื่อข้อมูลมีความยาวมากขึ้นซึ่งทำให้ไม่สามารถสังเกตุดูความเปลี่ยนแปลงของค่านี้ได้

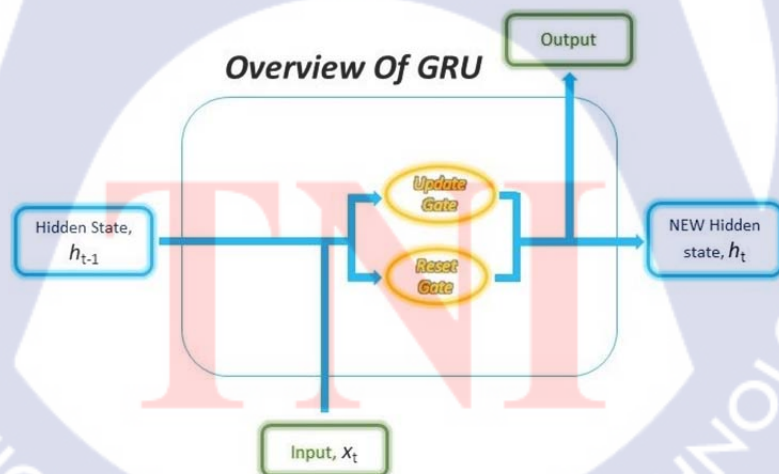
สิ่งที่ทำให้ LSTM แตกต่างจาก RNN นอกเหนือจากความเสถียรและประสิทธิภาพ ที่มากขึ้นคือ ความสามารถในการเก็บสถานะของแต่ละ Node ในโครงข่ายประสาทเทียมไว้เพื่อให้สามารถ ทราบค่าดั้งเดิมของ Node นั้นได้ ด้วยความสามารถนี้เองทำให้ LSTM สามารถกำหนดได้ถึงข้อมูลที่ควร จะจดจำและข้อมูลที่ควรจะถูกกำจัด



รูปที่ 2.7 Long Short-Term Memory

2.1.4.3 GRU

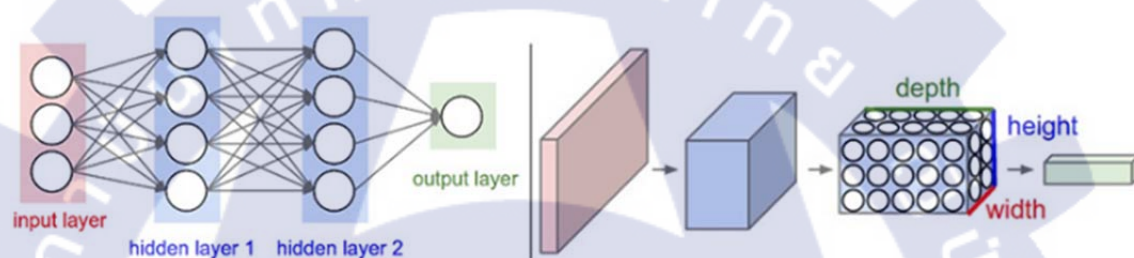
GRU ถูกพัฒนาต่อยอดมาจาก LSTM ซึ่งพัฒนาในส่วนของการลดความซับซ้อนในการทำงานของโครงข่ายประสาทแบบ LSTM โดย GRU ได้ทำการลดความซับซ้อนในการทำงานของโครงข่ายประสาทแบบ LSTM โดยการลดหน่วยย่อยใน Cell เหลือเพียง 2 ส่วน ได้แก่ Update Gate และ Reset Gate



รูปที่ 2.8 GRU

2.1.4.4 Convolutional Neural Network

Convolutional Neural Network หรือ CNN เป็นโครงข่ายประสาทเทียมแบบคอนโวลูชันแบบหนึ่งในกลุ่ม Bio-Inspired ซึ่งจะจำลองการมองเห็นของมนุษย์ที่มองเห็นพื้นที่เป็นที่ย่อยและมีการคัดแยกคุณลักษณะของพื้นที่ เช่น Contrast เป็นต้น โดย CNN ประกอบด้วยการคำนวณ Convolution หลายชั้นประกอบกับ Activation Functions เช่น ReLU ทำให้สามารถประมวลผลความสัมพันธ์ของข้อมูลแบบ Non-Linear ด้วยความสามารถนี้เองทำให้ CNN แตกต่างจากโครงข่ายประสาทเทียมแบบอื่นๆตรงที่สามารถดึงข้อมูลความสัมพันธ์เชิงพื้นที่ที่มีประโยชน์กับการประมวลผลข้อมูลรูปภาพได้การนำ CNN มาทำการประมวลผลภาษาธรรมชาตินั้นจะใช้ควบคู่กับการทำ Word Embedding โดยจะเป็นการนำ Word Vector ที่ได้จากการทำ Word Embedding มาทับซ้อนกันและประมวลผลเสมือนว่าเป็นข้อมูลรูปที่ 2.9



รูปที่ 2.9 Convolutional Neural Network

2.1.4.5 Word Embedding และ Word 2 Vector

Word Embedding คือการแปลงคำศัพท์เป็นตัวเลขให้อยู่ในรูปแบบของ Unit Vector โดย Word 2 Vector คือโมเดลที่ใช้ในการสร้าง Word Embedding ซึ่งจะใช้วิธีการคำนวณตัวเลขของคำนั้นๆ จากบริบทโดยรอบและการนำ Vector 2 ตัวมาทำการ Dot Product กันจะเป็นการหาค่าความคล้ายกันของ 2 Vector นี้ซึ่งเรียกได้ว่าเป็น Point-Wise Mutual Information (PMI) ดังนั้นจึงสามารถหาค่าความคล้ายกันของคำ (Word Similarity) ได้ด้วยวิธีนี้

การใช้ Word2Vec จะเป็นการสร้างโมเดลจากการทำ Word Embedding หลายชั้นมาสร้างเป็นโมเดล ซึ่งจะทำให้การฝึกโดยการคำนวณตัวเลขจากบริเวณใกล้เคียง (Context หรือ Corpus) กล่าวคือจะเป็นการคำนวณโดยอาศัยบริบทโดยรอบ

2.1.4.6 Keras

Keras คือ Programming Platform ที่พัฒนาเพื่อใช้กับภาษา Python ในการสร้างโมเดล โดยมีลักษณะเด่นคือใช้งานง่ายและรวดเร็วซึ่งต่อมาได้รับการสนับสนุนและพัฒนาจาก Google โดยตรงกลายเป็น API มาตรฐานของ Tensor Flow

2.1.4.7 Natural Language Toolkit

Natural Language Toolkit หรือ NLTK เป็นเครื่องมือที่ใช้ในการประมวลผลภาษาธรรมชาติที่ภายในประกอบไปด้วย Library ที่ใช้ในการประมวลผลข้อความโดยใช้ Apache License, Version 2.0 และรองรับทั้ง Python 2 และ Python 3

2.1.5 การตรวจสอบความถูกต้องของแบบจำลองการเรียนรู้เชิงลึก

Metrics เป็นตัววัดที่ใช้ในการประเมินความสามารถของโมเดลการเรียนรู้เชิงลึกว่าทำงานอย่างไรบ้าง มันช่วยให้เราวัดประสิทธิภาพของโมเดลได้อย่างเชื่อถือได้ และใช้ในการเปรียบเทียบระหว่างโมเดลที่แตกต่างกันได้ด้วย

นี่คือบาง Metrics ที่ใช้บ่อยในการตรวจสอบความถูกต้องของโมเดล

2.1.5.1 Confusion Matrix (เมทริกซ์ความสับสน) เป็นตารางที่แสดงจำนวนข้อมูลที่ถูกทำนายถูกต้องและผิดพลาด สามารถใช้ในการคำนวณ Metrics อื่นๆ ได้ เช่น Accuracy, Precision, Recall

2.1.5.2 Accuracy (ความแม่นยำ) จำนวนข้อมูลที่ถูกทำนายถูกต้องหารด้วยจำนวนข้อมูลทั้งหมดในชุดข้อมูลทดสอบ มันเหมาะสมสำหรับปัญหาที่คลาสที่สนใจมีความสมดุลกัน

2.1.5.3 Precision (ความแม่นยำในการทำนายบวก) Precision เป็นอัตราส่วนของตัวอย่างที่ถูกต้องที่โมเดลทำนายได้ (True Positive) ต่อทั้งหมดที่โมเดลทำนายว่าเป็นคลาสนั้นๆ (True Positive + False Positive)

$$\text{Precision} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}}$$

2.1.5.4 Recall (ความแม่นยำในการครอบคลุม) Recall เป็นอัตราส่วนของตัวอย่างที่ถูกต้องที่โมเดลทำนายได้ (True Positive) ต่อทั้งหมดของตัวอย่างที่เป็นคลาสนั้นๆ (True Positive + False Negative)

$$\text{Recall} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Negative}}$$

2.1.5.5 F1 Score

F1 score เป็นการวัดประสิทธิภาพของโมเดลในการจำแนกประเภทข้อมูล โดยคำนึงถึงความแม่นยำและความครอบคลุมของโมเดลทั้งสองด้าน ซึ่ง F1 Score นั้นคำนวณจากค่า Precision (ความแม่นยำ) และค่า Recall (ความครอบคลุม) ได้โดยใช้สูตรดังนี้

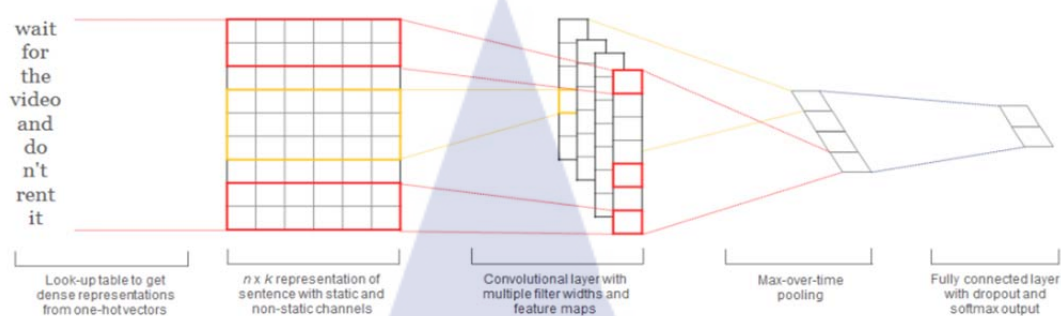
$$\text{F1 Score} = \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

F1 Score เป็นค่าที่อยู่ระหว่าง 0 ถึง 1 โดยค่ามากแสดงถึงประสิทธิภาพที่ดีของโมเดลในการจำแนกประเภท โดยมีความสมดุลระหว่างความแม่นยำและความครอบคลุม ซึ่งค่า F1 Score จะมีความเหมาะสมสำหรับการวัดประสิทธิภาพของโมเดลในกรณีที่ความสำคัญของ Precision และ Recall เท่าเทียมกัน

2.2 งานวิจัยที่เกี่ยวข้อง

แบบจำลอง Continuous Skip-Gram หรือ Word2Vec [15] เป็นวิธีการที่มีประสิทธิภาพสำหรับการเรียนรู้การแจกแจงของเวกเตอร์คุณภาพสูงที่รวบรวมความสัมพันธ์ของวากยสัมพันธ์และความหมายคำศัพท์ที่แม่นยำจำนวนมาก ในบทความนี้ได้นำเสนอส่วนขยายต่างๆ ที่ช่วยพัฒนาทั้งคุณภาพของเวกเตอร์และความเร็วในการฝึก โดยการสุ่มตัวอย่างคำที่ใช้บ่อย ผลลัพธ์ที่ได้มีความเร็วที่โดดเด่นและเรียนรู้ความหมายของคำต่างๆ ไปมากขึ้นด้วย แสดงให้เห็นว่าการเรียนรู้การแสดงเวกเตอร์ที่ดีสำหรับวลีนั้นเป็นไปได้

การทดลองการฝึกด้วย CNN ของ [16] แสดงให้เห็นว่าแบบจำลอง CNN ทัวไปเมื่อมีการปรับค่าไฮเปอร์พารามิเตอร์เล็กน้อย ก็จะได้ Static Vectors ที่ให้ผลลัพธ์ที่ยอดเยี่ยมในการวัดประสิทธิภาพ 7 รูปแบบ ดังรูปที่ 2.11

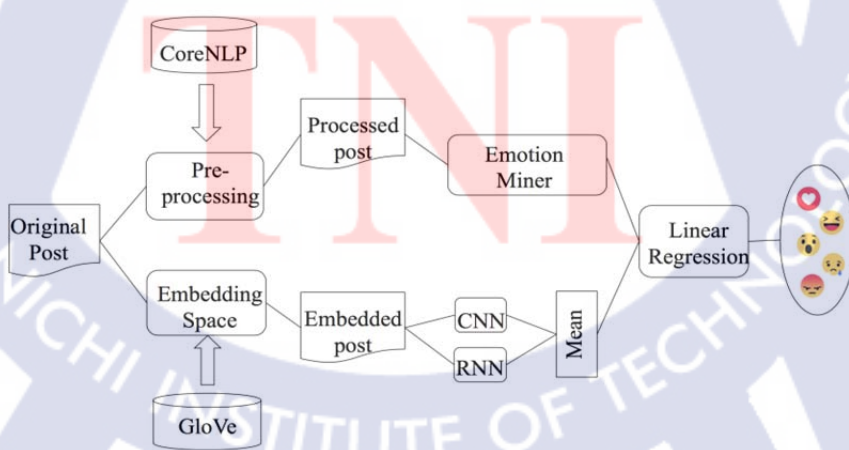


รูปที่ 2.10 Model CNN-Static [16]

Model	MR	SST-1	SST-2	Subj	TREC	CR	MPQA
CNN-rand	76.1	45.0	82.7	89.6	91.2	79.8	83.4
CNN-static	81.0	45.5	86.8	93.0	92.8	84.7	89.6
CNN-non-static	81.5	48.0	87.2	93.4	93.6	84.3	89.5
CNN-multichannel	81.1	47.4	88.1	93.2	92.2	85.0	89.4

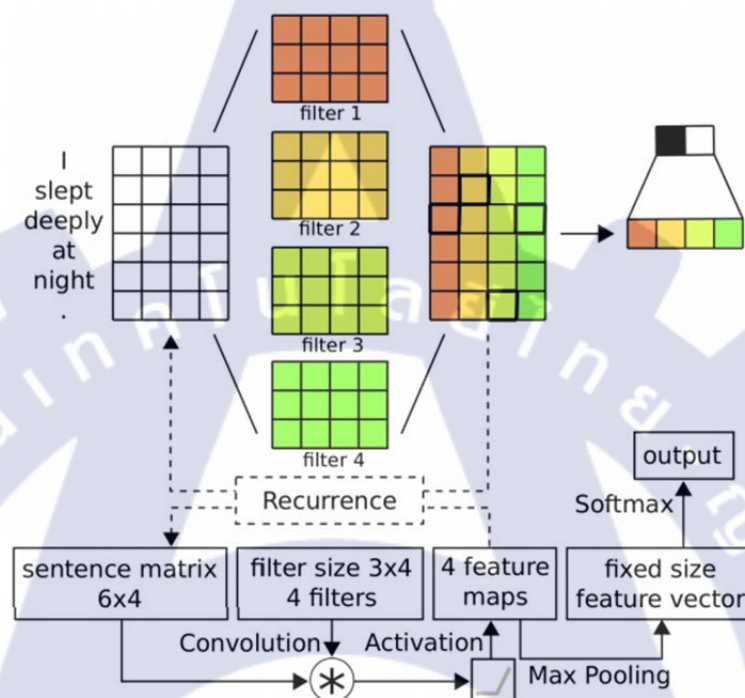
รูปที่ 2.11 ผลการทดลองแต่ละแบบจำลอง [16]

บนแพลตฟอร์ม Facebook มีฟังก์ชันที่สามารถแสดงปฏิกิริยาความรู้สึกต่อโพสต์ของผู้ใช้งาน ในหน้าสาธารณะของบริษัทต่างๆ [17] รวบรวมข้อมูลส่วนนี้ผ่าน Word Embedding ก่อนนำมาสร้างชุดข้อมูลสำหรับคาดการณ์ปฏิกิริยาโต้ตอบที่จะเกิดขึ้น โดยใช้ CNN และ RNN ร่วมกันในการฝึก และได้ผลลัพธ์ที่มีอัตราการทำนายผิดพลาดเพียง 0.135



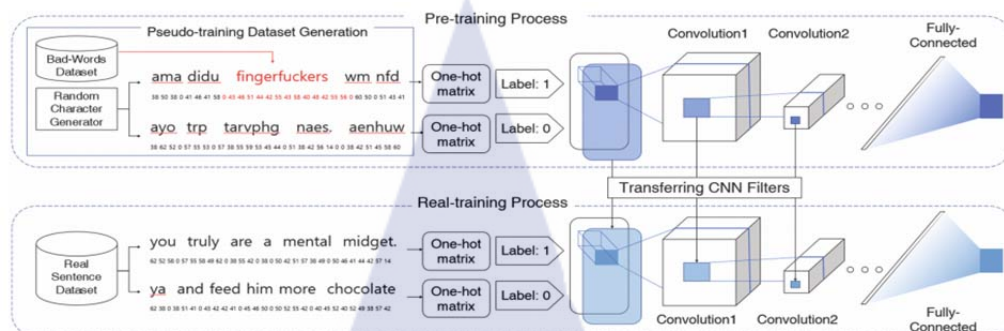
รูปที่ 2.12 Model Architecture ของแบบจำลอง [17]

การใช้แบบจำลอง CNN [18] ในการจำแนกประโยคในนวนิยาย บริบทในประโยคคือข้อมูลสำคัญสำหรับการทำความเข้าใจความหมาย แต่รูปแบบ Feed-Forward ของ CNN นั้นไม่เพียงพอที่จะสะท้อนถึงปัจจัยดังกล่าว เพื่อแก้ไขข้อจำกัดนี้ ใน Contextual CNN (C-CNN) จึงเพิ่ม RNN เข้ามาในแบบจำลองด้วย ทำให้แบบจำลองสามารถประมวลผลความหมายของคำร่วมกับสภาพแวดล้อมได้



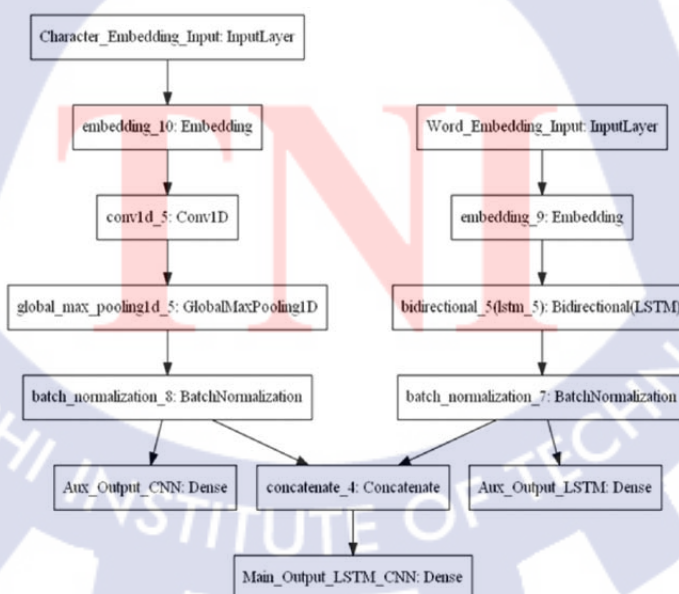
รูปที่ 2.13 Feed-Forward ของ CNN [18]

มีสองปัญหาในการจำแนกประโยคที่ไม่เหมาะสม หนึ่งคือการปรับเปลี่ยนของคำที่ไม่เหมาะสม และอีกอย่างคือระดับความไม่สมดุลซึ่งปรากฏในคลังข้อมูลเชิงรุกทั่วไป [19] ได้เสนอวิธีการฝึกประโยคปลอมที่สร้างเป็นระดับอักขระจนถึงการบิดเบือนเพื่อให้ข้อมูลไม่สมบูรณ์ เพื่อจัดการกับความไม่สมดุลของข้อมูล จึงเพิ่มคำศัพท์ที่มีความรุนแรงถึงครึ่งหนึ่งของประโยคที่สร้างขึ้นแบบสุ่ม เป็นชุดข้อมูลให้แบบจำลอง CNN ใช้ฝึกเมื่อนำ CNN ที่ผ่านการฝึกนี้แล้วไปใช้งานกับข้อมูลจริงก็พบว่ามีความแม่นยำมากขึ้น



รูปที่ 2.14 Model Architecture ของแบบจำลอง [19]

บนแพลตฟอร์มโซเชียลมีเดียอย่าง Twitter และ Facebook มีการใช้งานอย่างมหาศาลโดยองค์กรต่างๆ เพื่อเพิ่มช่องทางรับความคิดเห็นของผู้ใช้บริการเกี่ยวกับสถานการณ์ เหตุการณ์ ผลิตภัณฑ์และบริการ ดังนั้นการจัดประเภทความคิดเห็นจึงมีบทบาทสำคัญในการประเมินความคิดเห็นของผู้ใช้บริการ [20] ได้เสนอแบบจำลอง Novel Hybrid Deep Learning ที่รวมการ Word Embedding ต่างๆ (Word2Vec, Fast Text, Character-Level Embedding) ร่วมกับ Deep Learning ที่หลากหลาย (LSTM, GRU, BiLSTM, CNN) เพื่อตรวจสอบประสิทธิภาพของแบบจำลองดังกล่าว จึงมีการสร้างแบบจำลองหลายแบบเป็นแบบจำลองพื้นฐานสำหรับการทดลอง โดยเปรียบเทียบแต่ละแบบจำลองเพื่อหาวิธีการที่จะได้ผลลัพธ์ที่ดีขึ้น



รูปที่ 2.15 Layers ของแบบจำลอง M-Hybrid [20]

Model	Classification Method	Embedding	Accuracy (%)
M-1-A	CNN	Character-level	75.67
M-2	LSTM	Character-level	73.71
M-3	Bi-LSTM	Character-level	73.11
M-4	GRU	Character-level	74.57
M-5	CNN	FastText	76.71
M-6	LSTM	FastText	79.41
M-7-B	Bi-LSTM	FastText	80.44
M-8	GRU	FastText	79.91
M-9	CNN	Word2Vec	78.06
M-10	LSTM	Word2Vec	79.07
M-11	Bi-LSTM	Word2Vec	79.18
M-12	GRU	Word2Vec	79.59
M-Hybrid	CNN + BiLSTM	Character + FastText	82.14

รูปที่ 2.16 ผลทดลองความแม่นยำของแบบจำลองพื้นฐาน และ M-Hybrid [20]

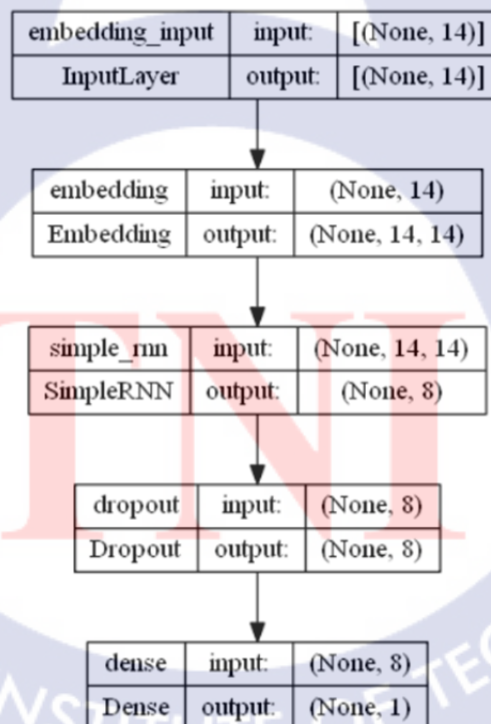
บทที่ 3

วิธีการวิจัย

3.1 กรอบความคิด

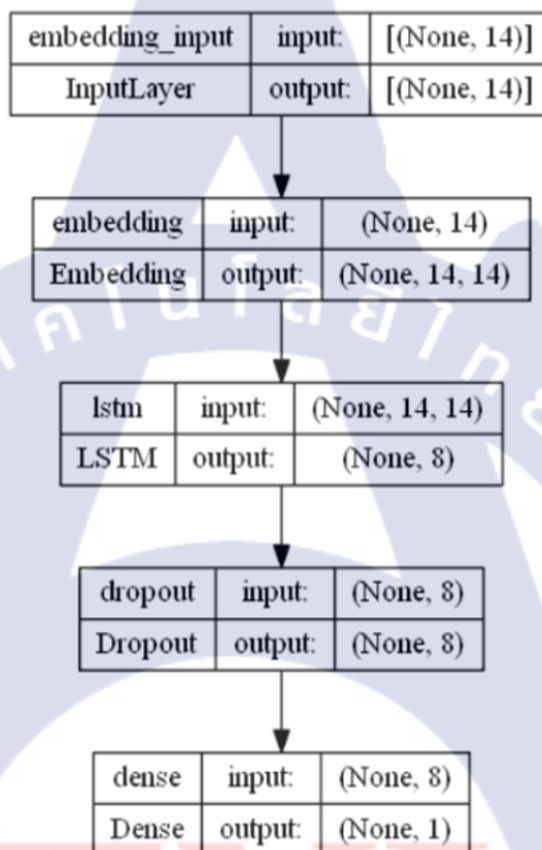
การวิเคราะห์ความรู้สึกของความคิดเห็นเชิงบวกและหรือเชิงลบเป็นกระบวนการดึงข้อมูลเกี่ยวกับการระบุตัวตนแบบอัตโนมัติจากข้อความที่สร้างขึ้นโดยผู้ใช้สื่อ ซึ่งอาจจะเป็นการจำแนกในระดับประโยค หรือในระดับเอกสาร และในปัจจุบันอาจจะใช้การวิเคราะห์ด้วยการเรียนรู้เชิงลึกโดยมีกระบวนการเพื่อเตรียมชุดข้อมูล เช่น TF-IDF, CBOW, Embedding, Skip-Gram ฯลฯ ดังที่กล่าวไว้ข้างต้น ดังนั้นเพื่อเพิ่มประสิทธิภาพและความเร็วในการประมวลผลต่อผู้ใช้งาน งานวิจัยนี้จึงนำเสนอระบบการวิเคราะห์ความรู้สึกโดยใช้ข้อมูลจากความคิดเห็นของผู้ใช้งานบนโซเชียลเน็ตเวิร์กโดยการนำการประมวลผลภาษาธรรมชาติ (NLP) ผสมกับอัลกอริทึมของโมเดลการเรียนรู้เชิงลึกเพื่อแก้ปัญหาการวิเคราะห์ความรู้สึกสำหรับการคัดแยกความรู้สึกเชิงบวก และเชิงลบ

อัลกอริทึมของโมเดลการเรียนรู้เชิงลึกที่ใช้ทั้งหมด มี 5 โมเดลดังรูปที่ 3.1-3.5



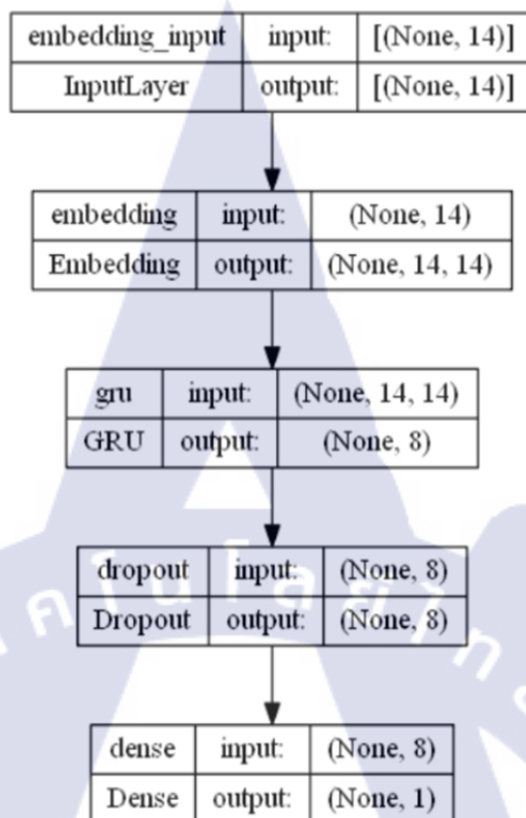
รูปที่ 3.1 RNN Model

โมเดลการเรียนรู้เชิงลึกที่ใช้ Recurrent Neural Network (RNN) มีประสิทธิภาพในการจำแนกประเภทและทำนายข้อมูลลำดับ เรื่องนี้เป็นไปได้ด้วยการใช้ลำดับข้อมูลในการอ้างอิงถึงสถานะก่อนหน้า และสร้างความเข้าใจเกี่ยวกับลำดับของเหตุการณ์ที่เกิดขึ้น อย่างไรก็ตาม RNN มีข้อจำกัดในการจดจำข้อมูลในช่วงที่ยาวนานเนื่องจากปัญหาการสูญเสยของข้อมูลในกระบวนการถดถอยที่ยาวนาน



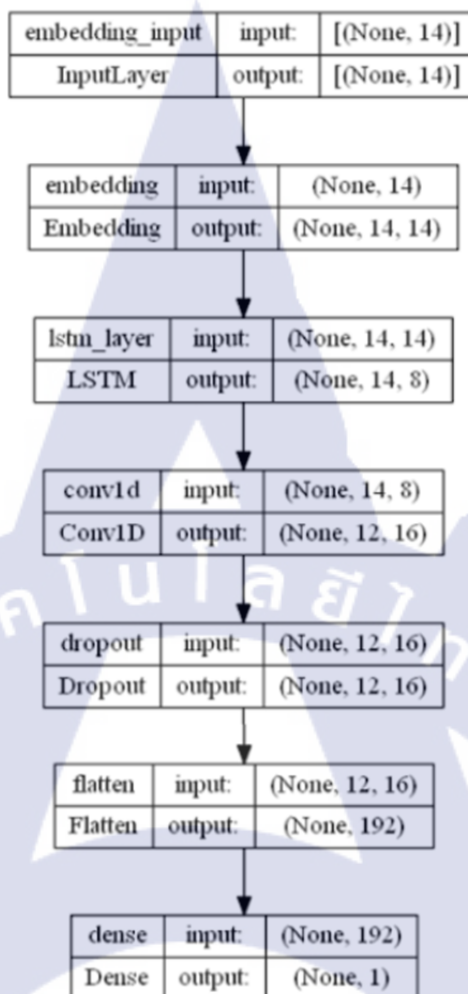
รูปที่ 3.2 LSTM Model

ประสิทธิภาพของ LSTM มาจากความสามารถในการจดจำข้อมูลในช่วงที่ยาวนานโดยไม่มีปัญหาการสูญเสยของข้อมูลในกระบวนการถดถอยที่ยาวนาน เป็นที่ประจักษ์ของปัญหาที่เจอในโมเดล RNN แบบพื้นฐาน



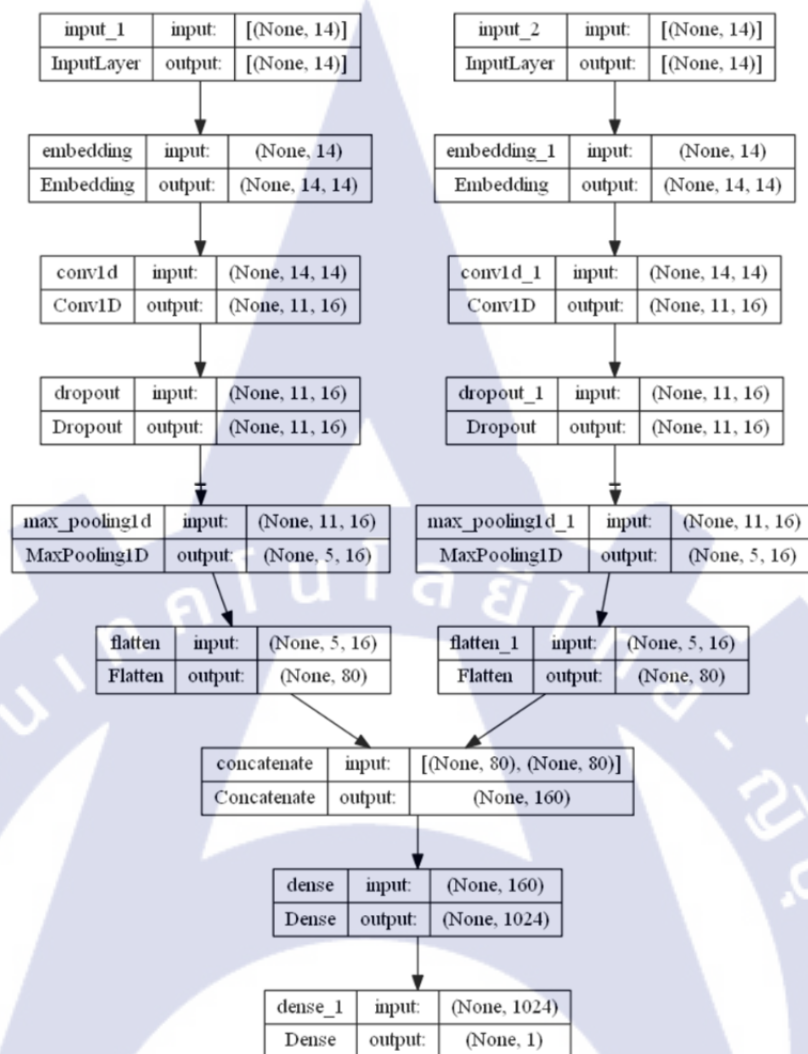
รูปที่ 3.3 GRU Model

โมเดลการเรียนรู้เชิงลึกที่ใช้ Gated Recurrent Unit (GRU) เป็นโมเดลที่สร้างมาจากแนวคิดของ LSTM โดยมีโครงสร้างที่เล็กกว่าและเร็วกว่า โมเดลนี้ถูกออกแบบมาเพื่อลดความซับซ้อนของ LSTM แต่ยังคงรักษาความสามารถในการจดจำข้อมูลในช่วงที่ยาวนานได้ดี



รูปที่ 3.4 CNN-LSTM Model

โมเดลการเรียนรู้เชิงลึกที่ใช้ CNN-LSTM เป็นการผสมผสานระหว่าง Convolutional Neural Network (CNN) และ Long Short-Term Memory (LSTM) ซึ่งเป็นสองโมเดลที่ได้รับความนิยมสูงในการจำแนกประเภทและทำนายข้อมูลลำดับในเดตตต่างๆ ให้กับ CNN เพื่อให้สามารถเข้าใจลักษณะทางพื้นที่ของข้อมูล และ LSTM เพื่อให้สามารถเข้าใจและจดจำลำดับของเหตุการณ์ได้ในข้อมูลลำดับ



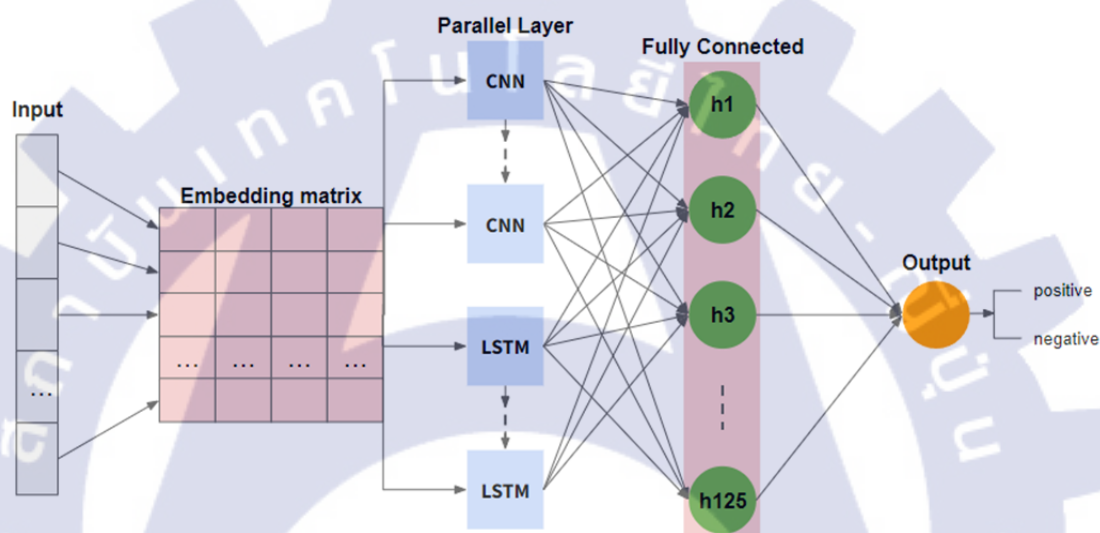
รูปที่ 3.5 Proposed Model

จากรูปที่ 3.5 อธิบายการทำงานของแบบจำลองจะรับข้อความอินพุตเดียวกันมาสองทาง นำมาผ่านกระบวนการ Embedding แปลงข้อความให้อยู่ในรูปแบบเวกเตอร์ตัวเลข เรียนรู้คุณลักษณะ จากข้อความโดยข้อความอินพุตที่ 1 ใช้ CNN และข้อความอินพุตที่ 2 ใช้ LSTM และนำเวกเตอร์ของข้อความ ที่ได้มารวมเข้าด้วยกันเพื่อส่งออกมาความน่าจะเป็นของข้อความอินพุต

ซึ่งโมเดลการเรียนรู้เชิงลึกดังกล่าวสามารถเรียนรู้รูปแบบที่ซับซ้อนในข้อมูลข้อความได้ดีกว่า วิธีการจำแนกประเภทแบบดั้งเดิม และสามารถจัดการกับข้อมูลขนาดใหญ่ได้

3.2 ขั้นตอนการทดลอง

ขั้นตอนวิธี (อัลกอริทึม) ของการพยากรณ์ความรู้สึกจากข้อความวิจารณ์ที่นำเสนอมีรายละเอียดดังต่อไปนี้ ประการแรกคือ ขั้นตอนของกระบวนการเตรียมข้อมูลซึ่งประกอบด้วยการทำความสะอาดข้อมูล (การแยกคำ การแยกประโยค การกรองตัวอักษรและลบอักขระต่างๆ และการปรับคำให้อยู่ในรูปรากของคำศัพท์นั้นๆ ตลอดจนการกำจัดคำศัพท์ที่ปรากฏน้อยออกไป) โดยหลังจากเตรียมชุดข้อมูลเป็นที่เรียบร้อยแล้ว ประการต่อมาคือการแปลงข้อความวิจารณ์เป็นเวกเตอร์ (หรือที่เรียกว่า Feature Vector) โดยใช้วิธี Word Embedding ซึ่งทำให้ได้ผลลัพธ์เวกเตอร์คุณลักษณะที่พร้อมจะนำไปป้อนให้กับอัลกอริทึมการเรียนรู้เชิงลึกที่ได้นำเสนอ ซึ่งในส่วนของโครงสร้างโมเดลที่นำเสนอแสดงได้ดังรูปที่ 3.6



รูปที่ 3.6 โครงสร้างของโมเดลการเรียนรู้เชิงลึกแบบผสมผสานที่นำเสนอ

ในส่วนนี้จะกล่าวถึงการทดสอบของระบบการคัดแยกความรู้สึกเชิงบวกและลบของข้อความคำวิจารณ์โดยมีขั้นตอนของการดำเนินงานดังรายละเอียดของแต่ละขั้นตอนดังนี้

การเตรียมชุดข้อมูลซึ่งประกอบด้วยชุดข้อมูลจากแหล่งต่างๆ โดยทั้งหมดนี้ ได้อธิบายไว้แล้วในบทที่ 2 หัวข้อที่ 2.1.1.1, 2.1.1.2 และ 2.1.1.3

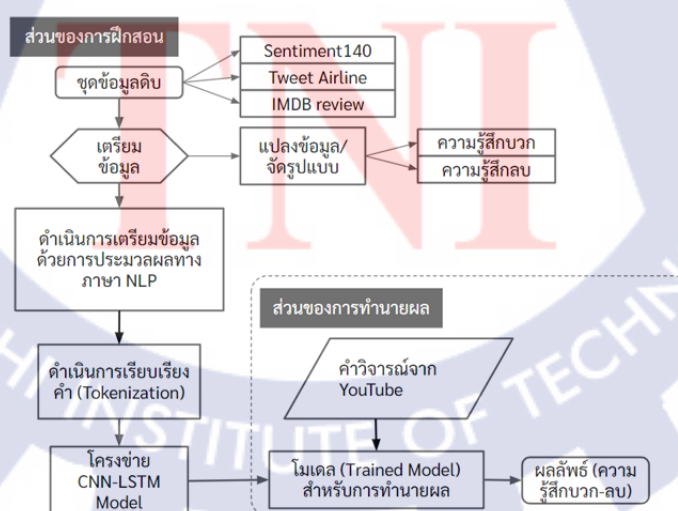
โดยทั้ง 3 ชุดข้อมูลสามารถแสดงตัวอย่างของข้อความวิจารณ์ดังรูปที่ 3.7

id	sentiment	review
5814_8	1	With all this stuff going down at the moment w...
2381_9	1	\The Classic War of the Worlds\" by Timothy Hi...
7759_3	0	The film starts with a manager (Nicholas Bell)...
3630_4	0	It must be assumed that those who praised this...
9495_8	1	Superbly trashy and wondrously unpretentious 8...
...	...	...
3471_3	0	this movie had more holes than a piece of swis...

รูปที่ 3.7 ตัวอย่างข้อมูลดิบสำหรับการทดลอง

หลังจากที่ได้ออกแบบอัลกอริทึมที่นำเสนอสำหรับการคัดแยกอารมณ์ความรู้สึกจากข้อความ วิจัยดังกล่าวไว้ในหัวข้อที่ผ่านมาแล้วนั้น ในขั้นตอนนี้เป็นดำเนินการฝึกสอนและทดสอบเพื่อ ประเมินประสิทธิภาพ การทดลองนี้ได้กำหนดค่าพารามิเตอร์ต่างๆ ที่สำคัญสำหรับการฝึกสอนได้แก่ รอบของการฝึกสอน (Epoch) เป็น 10 รอบ (ยกเว้นการทดลองกับชุดข้อมูลของ Tweets Airline และ Sentiment140 ของอัลกอริทึม GRU ตั้งค่าเป็น 15 รอบ) ชุดแบ่งข้อมูล (Batch Size) เท่ากับ 256, ค่าตัวแปร Loss = 'Binary_Crossentropy', และสุดท้ายตัวแปรของชนิดการเรียนรู้ (Optimizer) กำหนดให้ เป็น 'Adam'

3.3 การดำเนินการวิจัยและเก็บข้อมูล



รูปที่ 3.8 ขั้นตอนการดำเนินการวิจัย

3.4 ทำการทดลอง

3.4.1 การทดสอบการประมวลผลของแบบจำลอง

หาขั้นตอนการจำแนกที่ให้ผลลัพธ์ดีที่สุดจากการรวมสองขั้นตอนที่ให้ผลลัพธ์ที่ดีในด้านประสิทธิภาพและความแม่นยำ หลังจากที่ได้ดำเนินการฝึกสอนเป็นที่เรียบร้อยแล้วผลลัพธ์สามารถแสดงกราฟของการลดลงของค่าผิดพลาด (Loss) และค่าความถูกต้อง (Accuracy)

3.4.2 การทดสอบการประมวลผลด้วยชุดข้อมูลอื่น

ที่มาของชุดข้อมูลจะยังนำมาจากแพลตฟอร์ม Social Media Youtube โดยการจัดเก็บข้อความจากคอมเมนต์ของคลิปวิดีโอบนแพลตฟอร์ม ตัวอย่างข้อมูลตามรูปที่ 3.9

	review	sentiment
1	STREAM THIS ON SPOTIFY HERE!...	
2	sorry im late i was drooling looking at...	
3	NOW DO HOW TO EAT LIFE IT'...	
4	They should seriously make an anim...	
5	Gah demn nice voice bro	
6	What anime is it about	
7	I haven't stop listening to this ev...	
8	Что за аниме !?!?!?	
9	👍👍👍👍👍	
10	<a href="https://www.youtube.com/w...	

รูปที่ 3.9 ตัวอย่างข้อมูลดิบสำหรับการทดลอง

บทที่ 4

ผลการศึกษาและการอภิปรายผล

ในการดำเนินการทดลองและทดสอบทางผู้วิจัยได้ใช้ออกแบบอัลกอริทึมที่น่าเสนอ เพื่อฝึกสอนและทดสอบ(validate)ด้วยชุดข้อมูลที่เตรียมมาดังที่ได้กล่าวในหัวข้อ 3.2 ซึ่งได้แก่ ชุดข้อมูล Sentiment140 ชุดข้อมูล Tweets Airline และชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB โดยแต่ละการทดลองนั้นได้ดำเนินการใช้ชุดข้อมูลแต่ละแหล่งข้างต้นแยกออกจากกัน และแบ่งชุดข้อมูลสำหรับการฝึกสอนและทดสอบออกเป็นอัตราส่วน 80 ต่อ 20 ซึ่งสามารถแสดงผลการฝึกสอนได้ดังนี้

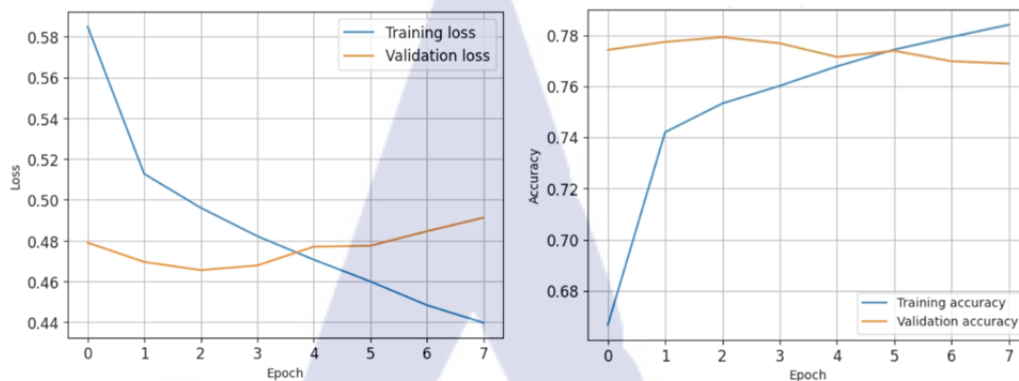
4.1 ผลการทดลอง

ตารางที่ 4.1 ร้อยละความถูกต้องของผลการทดลองการวิเคราะห์อารมณ์ความรู้สึกแบบและลบ

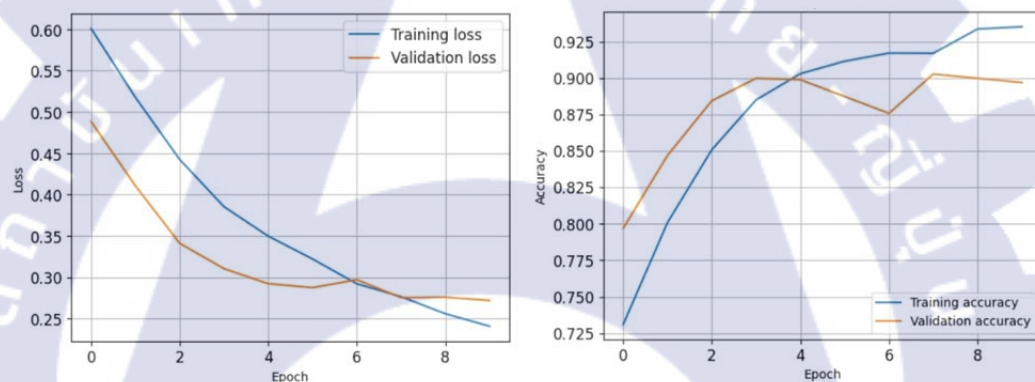
	ร้อยละความถูกต้องของผลการทดลองด้วยอัลกอริทึมการเรียนรู้เชิงลึกแบบต่างๆ				
	RNN	LSTM	GRU	CNN - LSTM	Proposed Model
Sentiment140 [12]	78.30%	78.48%	78.22%	77.40%	94.45%
IMDB Review [13]	85.16%	88.46%	87.46%	87.96%	95.67%
Tweet Airline [14]	91.08%	92.03%	91.52%	87.97%	97.41%

สำหรับการฝึกสอน จำนวนรอบของการฝึกสอนจะตั้งค่าที่ 10 รอบ เนื่องจากเมื่อทดสอบเกินกว่านั้นจะมีแนวโน้มทำค่าของความถูกต้องและค่าผิดพลาดของข้อมูลทดสอบ (Validation Data) นั้นสูงเกินไป หรือเกิดการ Overfitting กับข้อมูลที่ใช้ฝึกสอน ยกเว้นการฝึกสอนด้วยชุดข้อมูล Tweets Airline ของอัลกอริทึม GRU ที่ตั้งค่าจำนวนรอบของการฝึกสอนเป็น 15 รอบ และการฝึกสอนของอัลกอริทึม Proposed Model ที่ตั้งค่าจำนวนรอบของการฝึกสอนเป็น 50 รอบแทน และเพื่อให้เห็นถึงประสิทธิภาพของอัลกอริทึม ทางผู้วิจัยได้แสดงผลลัพธ์ของการทดลองที่อยู่ในรูปของ Confusion Matrix เพื่อบ่งชี้ว่าแต่ละอัลกอริทึมสำหรับการคัดแยกอารมณ์ความรู้สึกจากข้อความที่น่าเสนอนี้ มีความสามารถในการฝึกสอนให้สามารถคัดแยกความรู้สึกจากข้อความได้เป็นอย่างดี สำหรับการฝึกสอนของ ชุดข้อมูล Sentiment140 ชุดข้อมูล Tweets Airline และชุดข้อมูลบทวิจารณ์ภาพยนตร์ IMDB ตามลำดับ และในการทดลองนี้ยังได้เปรียบเทียบประสิทธิภาพของโมเดลกับโมเดลอื่นๆ ดังตารางที่ 4.2 ซึ่งจากผลที่แสดงในตารางนั้นจะเห็นว่าอัลกอริทึมที่น่าเสนอนี้มีความสามารถในการฝึกสอนให้สามารถคัดแยกความรู้สึกจากข้อความเฉลี่ยร้อยละ 95.84 จากตัวอย่างของการทดสอบทั้งหมด

4.1.1 RNN Model



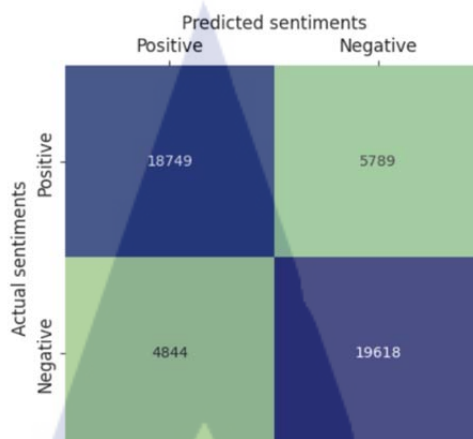
รูปที่ 4.1 กราฟของการทดลองของอัลกอริทึม RNN ด้วยชุดข้อมูล Sentiment140



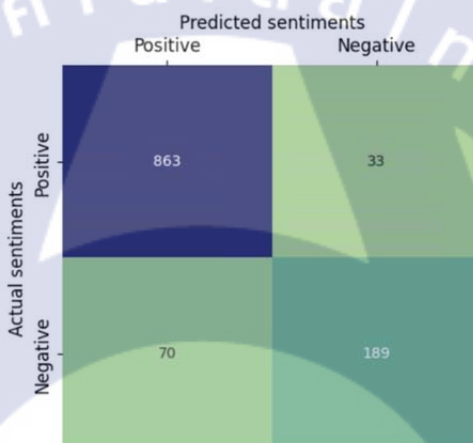
รูปที่ 4.2 กราฟของการทดลองของอัลกอริทึม RNN ด้วยชุดข้อมูล Tweets Airline



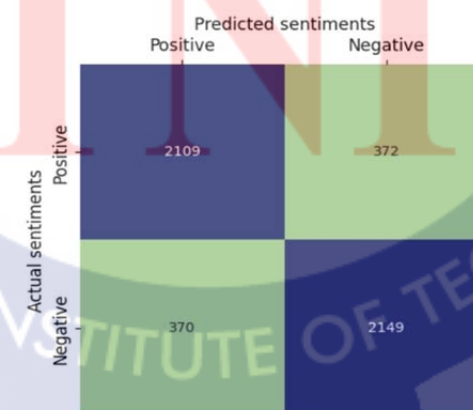
รูปที่ 4.3 กราฟของการทดลองของอัลกอริทึม RNN ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB



รูปที่ 4.4 ตาราง Confusion Matrix ของอัลกอริทึม RNN ด้วยชุดข้อมูล Sentiment140

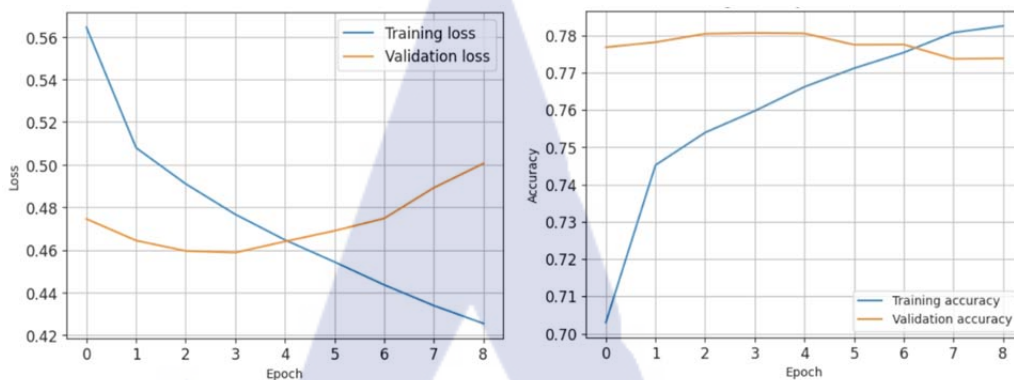


รูปที่ 4.5 ตาราง Confusion Matrix ของอัลกอริทึม RNN ด้วยชุดข้อมูล Tweets Airline

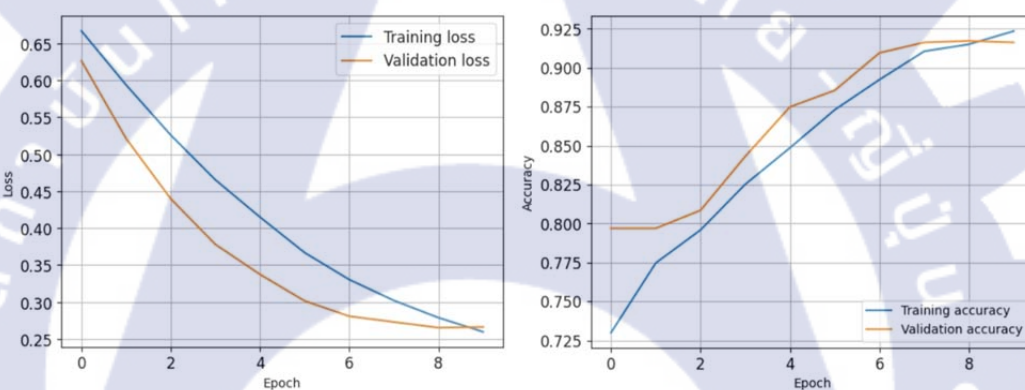


รูปที่ 4.6 ตาราง Confusion Matrix ของอัลกอริทึม RNN ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB

4.1.2 LSTM Model



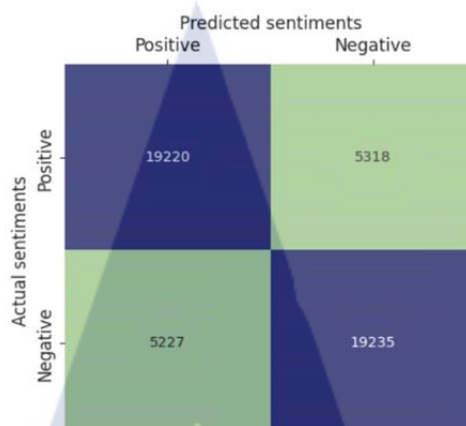
รูปที่ 4.7 กราฟของการทดลองของอัลกอริทึม LSTM ด้วยชุดข้อมูล Sentiment140



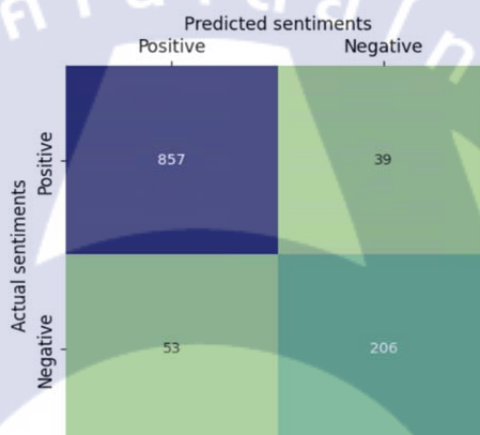
รูปที่ 4.8 กราฟของการทดลองของอัลกอริทึม LSTM ด้วยชุดข้อมูล Tweets Airline



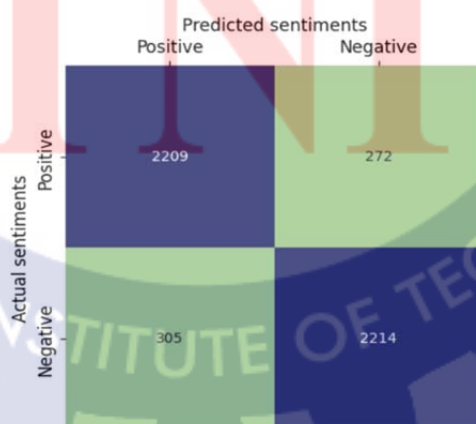
รูปที่ 4.9 กราฟของการทดลองของอัลกอริทึม LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB



รูปที่ 4.10 ตาราง Confusion Matrix ของอัลกอริทึม LSTM ด้วยชุดข้อมูล Sentiment140



รูปที่ 4.11 ตาราง Confusion Matrix ของอัลกอริทึม LSTM ด้วยชุดข้อมูล Tweets Airline

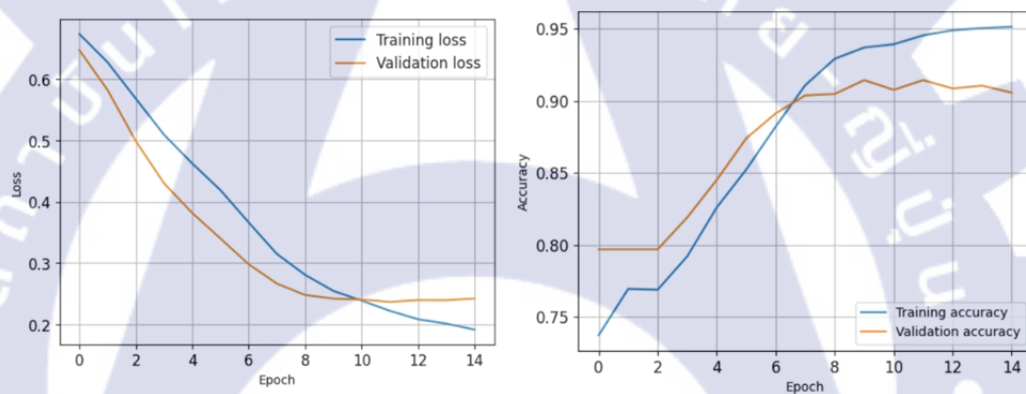


รูปที่ 4.12 ตาราง Confusion Matrix ของอัลกอริทึม LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB

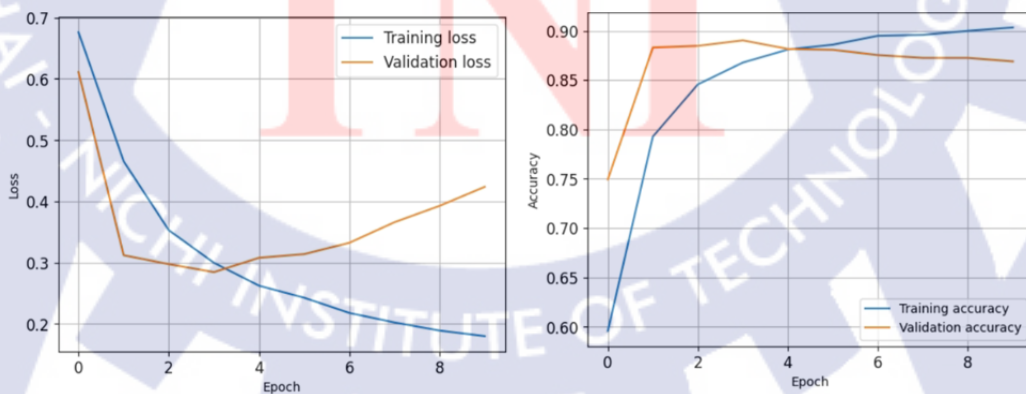
4.1.3 GRU Model



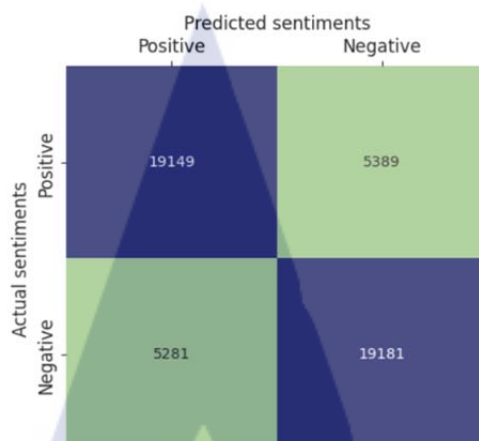
รูปที่ 4.13 กราฟของการทดลองของอัลกอริทึม GRU ด้วยชุดข้อมูล Sentiment140



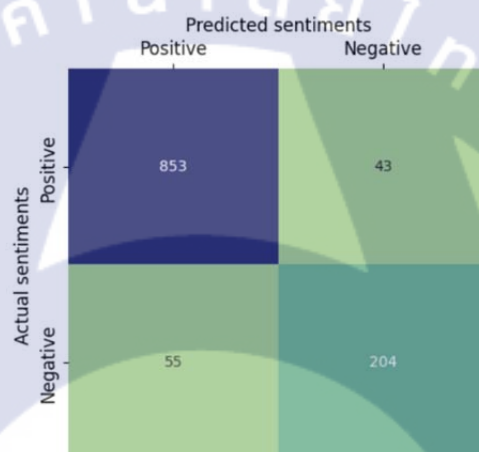
รูปที่ 4.14 กราฟของการทดลองของอัลกอริทึม GRU ด้วยชุดข้อมูล Tweets Airline



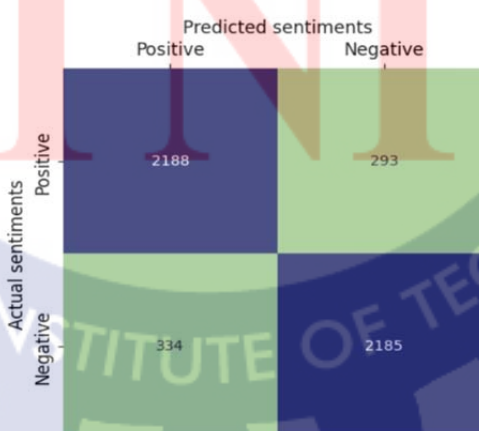
รูปที่ 4.15 กราฟของการทดลองของอัลกอริทึม GRU ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB



รูปที่ 4.16 ตาราง Confusion Matrix ของอัลกอริทึม GRU ด้วยชุดข้อมูล Sentiment140

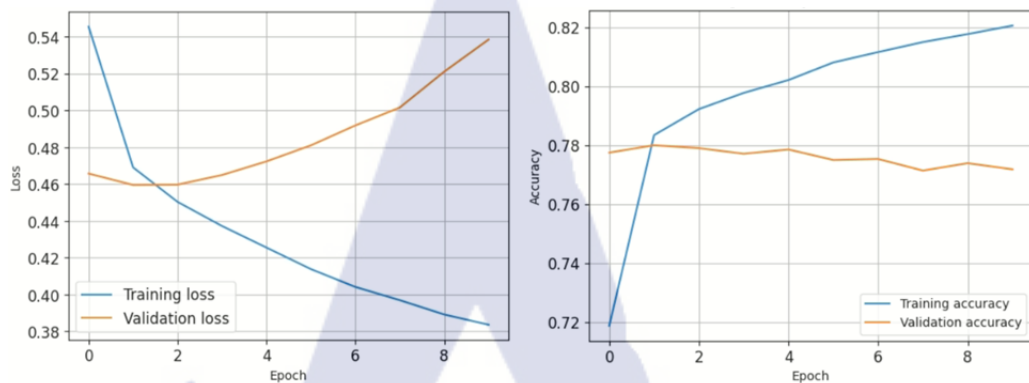


รูปที่ 4.17 ตาราง Confusion Matrix ของอัลกอริทึม GRU ด้วยชุดข้อมูล Tweets Airline

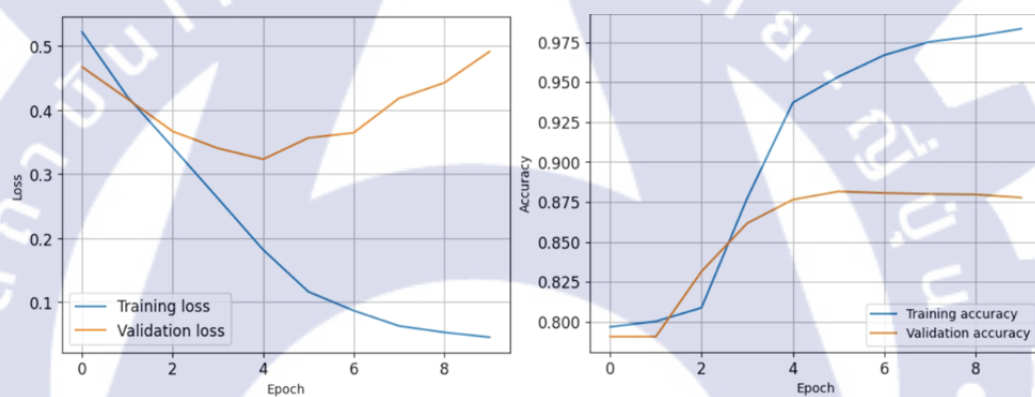


รูปที่ 4.18 ตาราง Confusion Matrix ของอัลกอริทึม GRU ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB

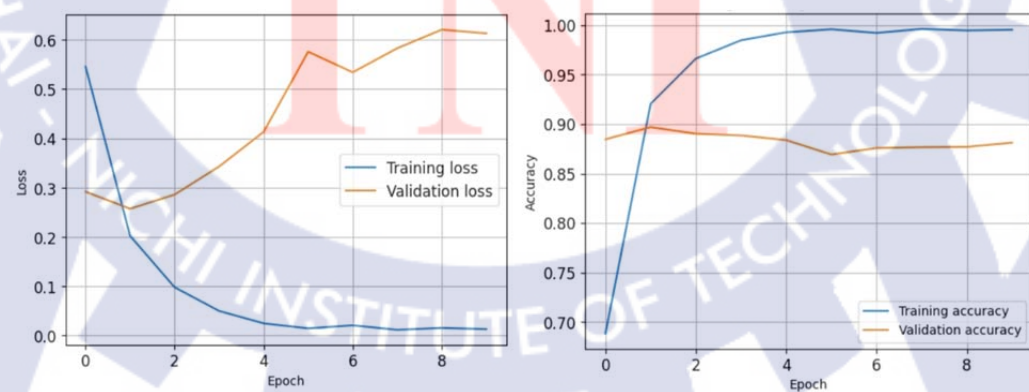
4.1.4 CNN-LSTM Model



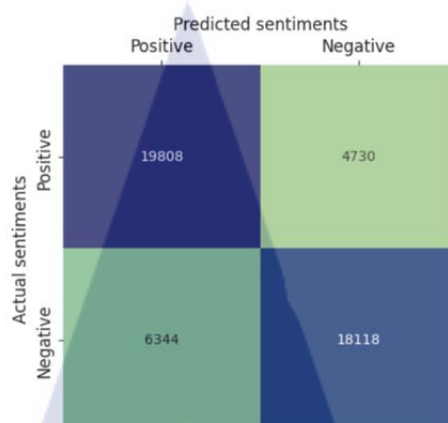
รูปที่ 4.19 กราฟของการทดลองของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Sentiment140



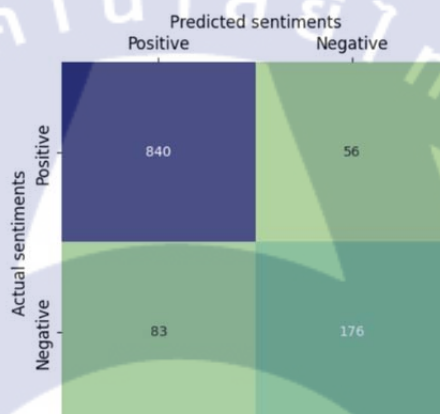
รูปที่ 4.20 กราฟของการทดลองของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Tweets Airline



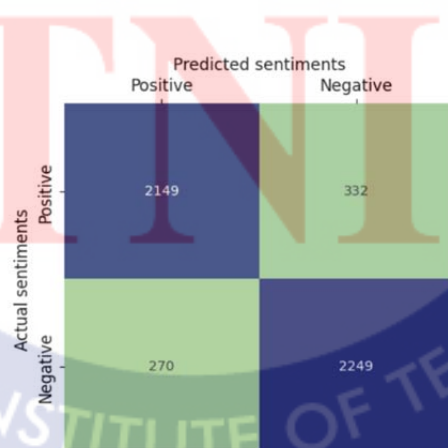
รูปที่ 4.21 กราฟของการทดลองของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB



รูปที่ 4.22 ตาราง Confusion Matrix ของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Sentiment140

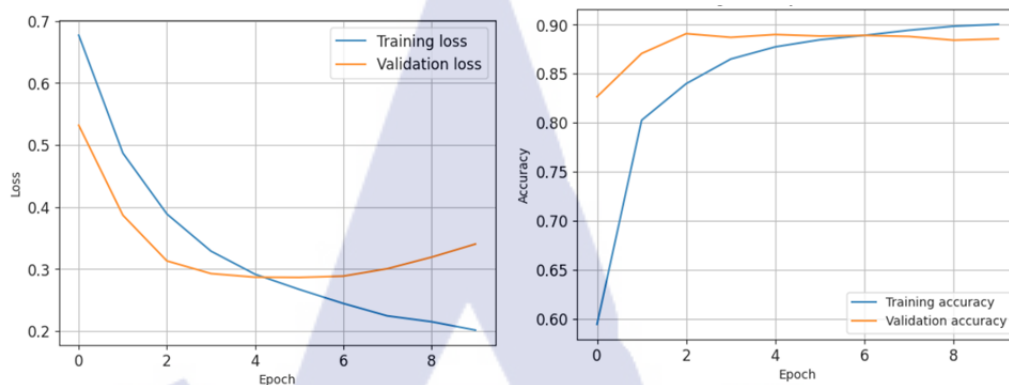


รูปที่ 4.23 ตาราง Confusion Matrix ของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล Tweets Airline

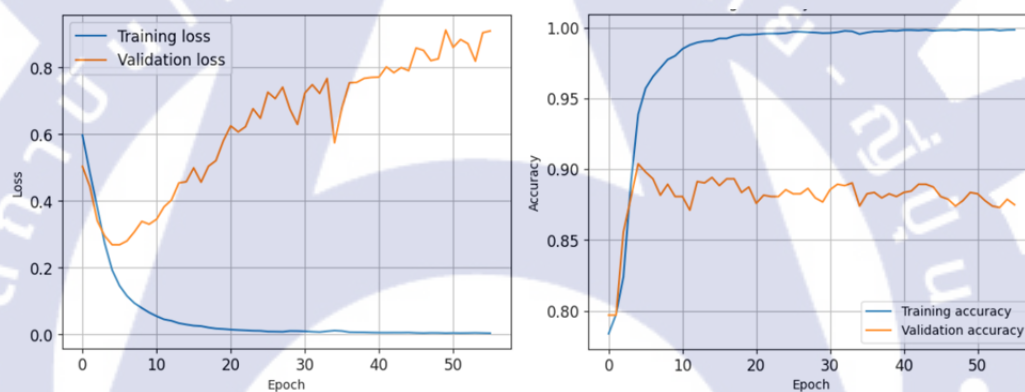


รูปที่ 4.24 ตาราง Confusion Matrix ของอัลกอริทึม CNN-LSTM ด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB

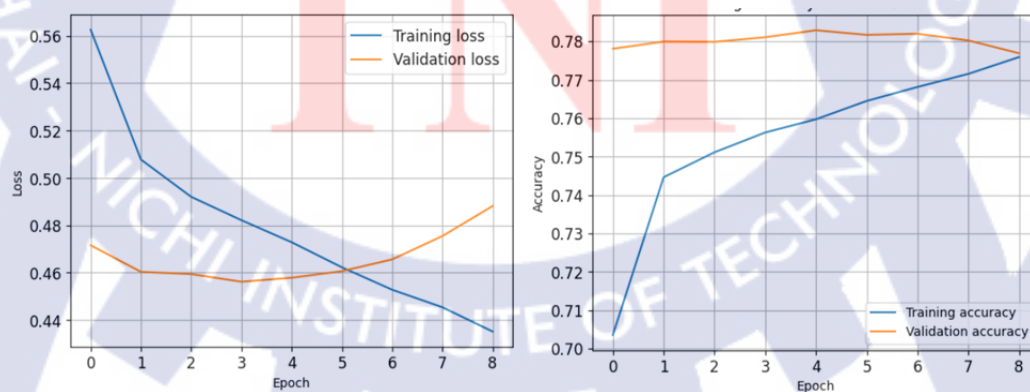
4.1.5 Propose Model



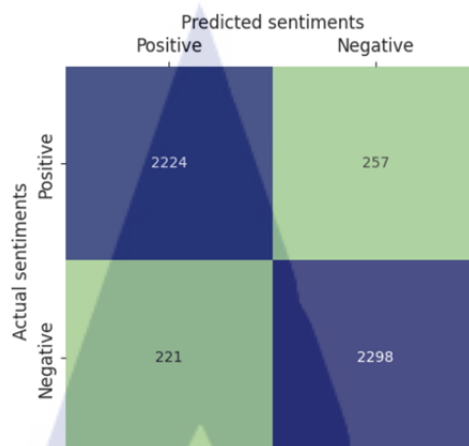
รูปที่ 4.25 กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วยชุดข้อมูล Sentiment140



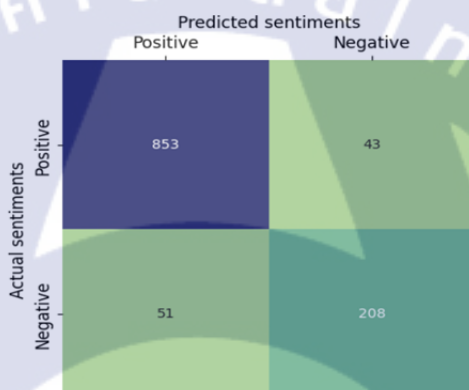
รูปที่ 4.26 กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วยชุดข้อมูล Tweets Airline



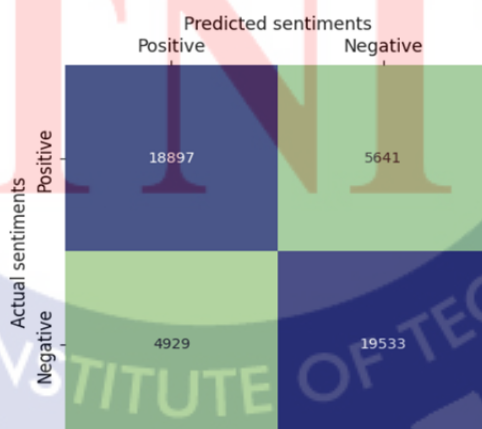
รูปที่ 4.27 กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วยชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB



รูปที่ 4.28 ตาราง Confusion Matrix นำเสนอด้วยชุดข้อมูล Sentiment140



รูปที่ 4.29 ตาราง Confusion Matrix นำเสนอด้วยชุดข้อมูล Tweets Airline



รูปที่ 4.30 ตาราง Confusion Matrix นำเสนอชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB

เมื่อนำตาราง Confusion Matrix จากการฝึกสอนแต่ละอัลกอริทึมมาคำนวณต่อตามสูตรการคำนวณ F1 Score เพื่อเปรียบเทียบประสิทธิภาพของโมเดลแต่ละโมเดล ได้ค่าออกมา ดังตารางที่ 4.2 ต่อไปนี้

ตารางที่ 4.2 ค่า F1 Score ของผลการทดลองการวิเคราะห์อารมณ์ความรู้สึกบวกและลบ

	F1 Score ของผลการทดลองด้วยอัลกอริทึมการเรียนรู้เชิงลึกแบบต่างๆ				
	RNN	LSTM	GRU	CNN - LSTM	Proposed Model
Sentiment140 [12]	0.7791	0.7847	0.7821	0.7815	0.9030
Tweet Airline [14]	0.9437	0.9491	0.9457	0.9236	0.7814
IMDB Review [13]	0.8504	0.8845	0.8747	0.8771	0.9487

บทที่ 5

สรุปผลและข้อเสนอแนะ

5.1 สรุปผล

งานวิจัยนี้ตระหนักเห็นถึงความจำเป็นอย่างยิ่งที่จะต้องพัฒนางานวิจัยในด้านการประมวลผลภาษาธรรมชาติ (NLP) โดยนำโมเดลการเรียนรู้เชิงลึกมาใช้เพื่อแก้ปัญหาคัดแยกความรู้สึกเชิงบวกและเชิงลบจากชุดข้อมูลที่เป็นข้อความวิจารณ์จากผู้บริโภคซึ่งอาจนำไปใช้ในการปรับปรุงผลิตภัณฑ์หรือสินค้าต่างๆ โดยในงานวิจัยนี้ได้ออกแบบขั้นตอนวิธี (อัลกอริทึม) ของการพยากรณ์ความรู้สึกจากข้อความวิจารณ์ที่ประกอบด้วยส่วนสำคัญอยู่ 2 ส่วนดังนี้ ซึ่งประกอบด้วย การทำความสะอาดข้อมูล (การแยกคำ การแยกประโยค การกรองตัวอักษรและลบอักขระต่างๆ และการปรับค่าให้อยู่ในรูปรากของคำศัพท์นั้นๆ ตลอดจนการกำจัดคำศัพท์ที่ปรากฏน้อยออกไป) โดยหลังจากเตรียมชุดข้อมูลเป็นที่เรียบร้อยแล้ว ประการต่อมาคือการแปลงข้อความวิจารณ์เป็นเวกเตอร์ (หรือที่เรียกว่า Feature Vector) โดยใช้วิธี Word Embedding ซึ่งทำให้ได้ผลลัพธ์เวกเตอร์คุณลักษณะที่พร้อมจะนำไปป้อนให้กับอัลกอริทึมการเรียนรู้เชิงลึกที่ได้นำเสนอ โดยผู้วิจัยได้ดำเนินการทดลองและทดสอบด้วยอัลกอริทึมที่นำเสนอเพื่อฝึกสอนและทดสอบด้วยชุดข้อมูล 3 ชุดได้แก่ ชุดข้อมูล Sentiment140 ชุดข้อมูล Tweets Airline และชุดข้อมูลบทวิจารณ์ภาพยนตร์ IMDB โดยแต่ละการทดลองนั้นได้ดำเนินการใช้ชุดข้อมูลแต่ละแหล่งข้างต้นแยกออกจากกัน และจากผลการทดลองจะได้ว่าอัลกอริทึมที่นำเสนอมีความสามารถในการฝึกสอนให้สามารถคัดแยกความรู้สึกจากข้อความเฉลี่ยในระดับที่ดีมาก สามารถคัดแยกความรู้สึกจากข้อความเฉลี่ยร้อยละ 95.84 จากตัวอย่างของการทดสอบทั้งหมด

5.2 ข้อเสนอแนะ

5.2.1 การจำลองนี้ใช้ชุดข้อมูลมาตรฐาน 3 ชุดในการเรียนรู้ ซึ่งชุดข้อมูลดังกล่าวอาจมีข้อมูลและหลักการใช้ภาษาที่ไม่ตรงกับค่านิยมในปัจจุบัน หากต้องการนำไปใช้ ควรรวบรวมข้อมูลคำวิจารณ์ที่เกี่ยวข้องเพิ่มสำหรับการเรียนรู้ของโมเดล

5.2.2 การใช้โมเดลร่วมกันมากกว่า 1 โมเดล ช่วยเพิ่มความแม่นยำ ความถูกต้อง หรือลดการใช้เวลาประมวลผลเพิ่มขึ้นได้ ขึ้นอยู่กับความสามารถของแต่ละอัลกอริทึมหรือการออกแบบโมเดล

5.2.3 โมเดลอัลกอริทึมแบบคู่ขนานและแบบตามลำดับมีหลากหลายรูปแบบ ซึ่งการทดลองในครั้งนี้เป็นเพียงหนึ่งในการออกแบบโมเดลอัลกอริทึมที่ช่วยเพิ่มประสิทธิภาพในการประมวลผล

5.2.4 โมเดลการเรียนรู้เชิงลึกต้องการข้อมูลจำนวนมากเพื่อฝึกให้มีประสิทธิภาพ หากไม่มีข้อมูลเพียงพอ โมเดลอาจมีประสิทธิภาพต่ำหรือเกิดการ Overfitting กับชุดข้อมูล Training ได้

5.2.5 โมเดลการเรียนรู้เชิงลึกอาจสะท้อนอคติในข้อมูลที่ใช้ฝึก ซึ่งอาจนำไปสู่ผลลัพธ์ที่ไม่เป็นธรรมหรือไม่ถูกต้อง

5.2.6 ควรปรับแต่งโมเดลเพิ่มเติมเพื่อป้องกัน Overfit และเพิ่มประสิทธิภาพวิเคราะห์บนชุดข้อมูลที่ไม่เคยเห็นมาก่อน



บรรณานุกรม

TNI

บรรณานุกรม

- [1] M. Lukauskas et al., "Economic activity forecasting based on the sentiment analysis of news," *Mathematics*, vol. 10, no. 19, pp. 3,461, September 2022.
- [2] A. J. Vikram, "Sentiment analysis and accuracy prediction of medical dataset using deep adversarial networks," *Quing Publications*, vol. 1, no. 2, pp. 76-78, April-June 2022.
- [3] H. Grissette and E. H. Nfaoui, "Deep associative learning approach for bio-medical sentiment analysis utilizing unsupervised representation from large-scale patients' narratives," *Pers Ubiquitous Compute*, vol. 27, no. 3, pp. 2,055-2,069, August 2021.
- [4] K. Mishev et al., "Evaluation of sentiment analysis in finance : From lexicons to transformers," *KKU Engineering Journal*, vol. 8, no. 5, pp. 131-136, July 2020.
- [5] E. Kauffmann et al., "Managing marketing decision-making with sentiment analysis : An evaluation of the main product features using text data mining," *Journal Teleoperators and Virtual Environments*, vol. 11, no. 15, pp. 4.235, August 2019.
- [6] M. G. Sousa et al., "Bert-based stock market sentiment analysis," *IEEE International Conference on Tools with Artificial Intelligence, ICTAI 2019, Portland, USA, November 4-6, 2019*, pp. 1,597-1,601.
- [7] Z. Li et al., "Multimodal sentiment analysis based on interactive transformer and soft mapping," *Advances in Neural Networks and Applications Journal*, vol. 10, no. 1, pp. 287-291, February 2022.
- [8] M. Hadwan et al., "An improved sentiment classification approach for measuring user satisfaction toward governmental services' mobile apps using machine learning methods with feature engineering and smote technique," *Applied Sciences*, vol. 12, no. 11, pp. 5,547, May 2022.
- [9] M. R. Qureshi et al., "A simple approach to classify fictional and non-fictional genres," *International Conference on Information and Communication Technology Convergence, ICIT 2019, Florence, Italy, August 1, 2019*, pp. 81-89.
- [10] M. U. Salur and I. Aydin, "A novel hybrid deep learning model for sentiment classification," *Indian Chemical Engineer Journal*, vol. 8, no. 1, pp. 37-48, March 2020.

- 
- [11] C. N. Dang et al., “Hybrid deep learning models for sentiment analysis,” *Agriculture and Human Journal*, vol. 31, no. 1, pp. 33-51, August 2021.
 - [12] A. Go et al., “*Sentiment140*,” [Online]. Available : <http://help.sentiment140.com/site-functionality>. [Accessed : February 12, 2023].
 - [13] F. Eight and K. Team “*Twitter US Airline Sentiment*,” [Online]. Available : <https://www.kaggle.com/crowdflower/twitter-airline-sentiment>. [Accessed : February 12, 2023].
 - [14] A. Maas “*IMDB Dataset of 50K Movie Reviews*,” [Online]. Available : <https://www.kaggle.com/datasets/lakshmi25npathi/imdb-dataset-of-50k-movie-reviews>. [Accessed : February 12, 2023].
 - [15] T. Mikolov et al., “Distributed representations of words and phrases and their compositionality,” *Neural Information Processing Systems*, ICIT 2013, New York, USA, December 5-10, 2013, pp. 3,111-3,119.
 - [16] Y. Kim, “Convolutional neural networks for sentence classification,” *Empirical Methods in Natural Language Processing*, EMNLP 2014, Doha, Qatar, October 25-29, 2014, pp. 1,746-1,751.
 - [17] F. Krebs et al., “Social emotion mining techniques for facebook posts reaction prediction,” *International Conference on Agents and Artificial Intelligence*, ICAART 2018, Funchal, Madeira, Portugal, January 16-18, 2018, pp. 211-220.
 - [18] J. Shin et al., “Contextual-CNN : A novel architecture capturing unified meaning for sentence classification,” *IEEE International Conference on Big Data and Smart Computing*, Big Comp 2018, Shanghai, China, January 15-17, 2018, pp. 491-494.
 - [19] S. Seo and S. B. Cho, “Offensive sentence classification using character-level CNN and transfer learning with fake sentences,” *International Conference on Neural Information Processing*, ICONIP 2017, Guangzhou, China, November 14-18, 2017, pp. 532-539.



ภาคผนวก

ภาคผนวก ก.
บทความการประชุมวิชาการ



การวิเคราะห์ความคิดเห็นจากข้อความความคิดเห็นด้วยการเรียนรู้เชิงลึก

Sentiment Classification from Text Comments using a hybrid CNN-LSTM

นางสาวคณิน พงศ์ศักดิ์ศรี*, ดร. อัดนา เซนต์เดะ*

Kaninard Pongsaksr, Dr. Adna Sento

คณะวิศวกรรมศาสตร์, สถาบันเทคโนโลยีไทย-ญี่ปุ่น, ประเทศไทย

Corresponding Author E-mail addresses: po.kaninard_si@tni.ac.th, adna@tni.ac.th

บทคัดย่อ — ในปัจจุบันการวิเคราะห์ข้อมูลจากความคิดเห็นของประชาชนบนโซเชียลเน็ตเวิร์ก เช่น YouTube หรือ Twitter หรืออื่นๆ ได้กลายมามีบทบาทสำคัญสำหรับการนำไปใช้งานด้านต่างๆ เช่น การตลาด การเมือง ฯลฯ โดยเฉพาะอย่างยิ่งงานด้านการพยากรณ์อารมณ์ความรู้สึกเชิงบวกและเชิงลบจากความคิดเห็นของผู้ที่ใช้งานโซเชียลเน็ตเวิร์กเพื่อใช้ในการปรับปรุงผลิตภัณฑ์หรือสินค้าต่างๆ นั้นเป็นสิ่งที่ต้องการอย่างมาก ดังนั้นงานวิจัยนี้จึงนำเสนอโมเดลการเรียนรู้เชิงลึกแบบผสมผสานกับการประมวลผลภาษาธรรมชาติ (NLP) สำหรับการคัดแยกความรู้สึกเชิงบวกและเชิงลบ โดยในงานวิจัยนี้ได้ออกแบบขั้นตอนวิธี (อัลกอริทึม) ของการพยากรณ์ความรู้สึกจากข้อความวิจารณ์ที่ประกอบด้วยส่วนสำคัญอยู่ 2 ส่วนดังนี้ ส่วนแรกเป็นขั้นตอนของกระบวนการเตรียมข้อมูลซึ่งประกอบด้วยการทำความสะอาดข้อมูลและการดึงคุณลักษณะเด่นของข้อมูล และส่วนที่สองคือการออกแบบโมเดลสำหรับการเรียนรู้ด้วยการใช้โมเดลการเรียนรู้เชิงลึกแบบผสมผสาน และในตอนท้ายของงานวิจัยนี้ก็ยังได้ออกแบบการทดลองเพื่อพิสูจน์ประสิทธิภาพของโมเดลอัลกอริทึมที่ได้นำเสนอ ซึ่งผลลัพธ์ของการทดลองก็สามารถคัดแยกความรู้สึกจากข้อความเฉลี่ยร้อยละ 95.84 จากตัวอย่างของการทดสอบทั้งหมด

คำสำคัญ — การวิเคราะห์ความรู้สึก, การเรียนรู้เชิงลึก, การเรียนรู้ของเครื่อง, โครงข่ายประสาทเทียม, การประมวลผลภาษาธรรมชาติ

ABSTRACT — Analyzing data from public opinion on social networks such as YouTube, Twitter, etc., has become essential for applications in fields such as marketing, politics, etc. in particular, the task of forecasting the positive and the negative emotions from the user opinions on social networks to improve products or products has become highly in demand. Therefore, this research presents a deep learning model combined with natural language processing (NLP) for the classification of the positive and negative sentiments. In this research, an algorithm (algorithm) for sentiment classification from commentary text was designed that consists of two main parts as follows: the first part is the data preparation process which includes data cleaning and extraction of the data features. The second part is designing a model for learning using the integrated deep learning model. At the end of this research, an experiment was designed to prove the efficiency of the proposed algorithmic model. The experiment's results extracted feelings from messages on average 95.84 percent from all samples of the tests.

Keywords — sentiment analysis, deep learning, machine learning, neural network, natural language processing.

1. บทนำ

จาก [1] ได้ให้ข้อมูลเกี่ยวกับพฤติกรรมของผู้บริโภคในยุคปัจจุบันที่ใช้ช่องทางการสั่งซื้อออนไลน์มากขึ้น จุดนี้เองทำให้เกิดการวิจัยสินค้าในกลุ่มผู้บริโภคต่างๆ ที่เขียนลงบนโซเชียลมีเดียมากขึ้นตามไปด้วย ด้วยเหตุนี้ทำให้เกิดการตัดสินใจของผู้บริโภคต่อสินค้าใดๆ ก่อนการเลือกซื้อด้วยการอ่านคำวิจารณ์ของสินค้าเสมอ ในทางกลับกันผู้ผลิตสินค้าก็สามารถพัฒนาสินค้าของตนด้วยการนำคำวิจารณ์เหล่านั้นมาปรับปรุงและพัฒนาสินค้าของตัวเองอีกด้วย และปัจจุบันก็เป็นที่ยอมรับกันอย่างกว้างขวางว่าการวิเคราะห์ความรู้สึกนั้นมีประโยชน์อย่างมากในการดำเนินการด้านต่างๆ เช่น ระบบธุรกิจอีคอมเมิร์ซและอีคอมเมิร์ซ ซึ่งทำให้สามารถศึกษาความคิดเห็นของลูกค้าเพื่อทำการสนับสนุนลูกค้าที่ดีขึ้น สร้างผลิตภัณฑ์ที่ดีขึ้น หรือปรับปรุงกลยุทธ์ทางการตลาดเพื่อดึงดูดลูกค้าใหม่ สามารถใช้การวิเคราะห์ความรู้สึกเพื่อสรุปความคิดเห็นของผู้ใช้เกี่ยวกับเหตุการณ์หรือผลิตภัณฑ์ ผลลัพธ์ของการวิเคราะห์สามารถช่วยให้ได้รับข้อมูลเชิงลึกมากขึ้นเกี่ยวกับความสนใจหรือความคิดเห็นของลูกค้าเกี่ยวกับแนวโน้มอุตสาหกรรม ส่งผลให้ทีมงานวิจัยที่มุ่งเน้นเกี่ยวกับงานการสร้างแบบจำลองคุณภาพสูงเพื่อแก้ปัญหาความซับซ้อนของข้อมูลขนาดใหญ่มีแนวโน้มที่เพิ่มขึ้นอย่างต่อเนื่อง เพื่อการพัฒนาทางด้านงานด้านการนำผลจากการศึกษาไปใช้ในส่วนของการขยายการวิเคราะห์ความรู้สึกไปสู่การใช้งานที่หลากหลายไม่ว่าจะเป็นด้านสุขภาพ [2,3] ด้านการเงิน [4] ด้านการวิเคราะห์การตลาด [5,6] และอื่นๆ [7,8,9,10] ยิ่งไปกว่านั้นเมื่อไม่นานมานี้การนำแนวคิดนี้สร้างงานวิจัยเกี่ยวกับการวิเคราะห์พฤติกรรมของผู้บริโภคจากความคิดเห็นหรือความวิจารณ์แบบอัตโนมัติมากขึ้นเรื่อยๆ เช่นกัน ตัวอย่างเช่น [11]

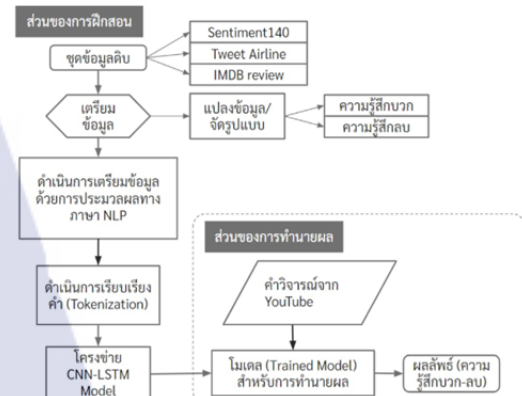
โดยมีโมเดลของการเรียนรู้เชิงลึกเป็นตัวขับเคลื่อนให้ทีมงานวิจัยที่มีประสิทธิภาพมากขึ้นซึ่งส่วนใหญ่ของงานวิจัยได้นำ Convolutional Neural Networks (CNN) หรือ Recurrent Neural Networks (RNN) หรือ Long Short-Term Memory (LSTM) ที่ให้ผลลัพธ์ความแม่นยำโดยรวมที่ค่อนข้างสูงเพื่อให้งานวิจัยได้รับการยอมรับมากขึ้น ด้วยเหตุผลนี้เองในงานวิจัยนี้จึงตระหนักถึงความจำเป็นอย่างยิ่งที่จะต้องพัฒนางานวิจัยในด้านนี้เพื่อเพิ่มและปรับปรุงประสิทธิภาพในการเรียนรู้ความคิดเห็นของประชากรผู้บริโภคในลักษณะการสั่งซื้อออนไลน์ หรือบนพื้นที่โซเชียลโดยกระบวนการของงานวิจัยนี้ได้้นำการประมวลผลภาษาธรรมชาติ (NLP) มาผนวกกับอัลกอริทึมของโมเดลการเรียนรู้เชิงลึกเพื่อแก้ปัญหาการ

ต้องการ โดยทั่วไป Word Embedding สามารถสร้างขึ้นได้โดยใช้วิธีการต่างๆ เช่น โครงข่ายประสาทเทียม เมทริกซ์แบบพิเศษ แบบจำลองความน่าจะเป็น ฯลฯ Word2Vec สามารถทำได้ในรูปแบบ Continuous Bag of Word (CBOW) และ Skip-gram โดย CBOW เป็นโมเดลที่ทำนายคำปัจจุบันที่กำหนดคำบริบทภายในลำดับคำอื่นๆ โครงข่ายของ CBOW ประกอบด้วยอินพุตที่เป็นคำที่อยู่ลำดับใกล้ๆ กัน และเอาต์พุตกำหนดให้เป็นคำปัจจุบัน ส่วนชั้นซ่อน (Hidden Layer) เป็นตัวกำหนดมิติที่เราต้องการแสดงคำปัจจุบัน สำหรับ Skip-gram จะเป็นโมเดลที่ทำนายคำโดยรอบ โดยโครงสร้างเลเยอร์ของ Skip-gram ประกอบด้วยอินพุต (คำปัจจุบัน) และเลเยอร์เอาต์พุตกำหนดให้เป็นคำโดยรอบ ส่วนชั้นซ่อน (Hidden Layer) เป็นตัวกำหนดมิติที่เราต้องการแสดงคำปัจจุบัน

2.3 โครงข่ายประสาทเทียม (Neural Network) และส่วนของเอาต์พุต
ในส่วนของการคัดแยกกลุ่มข้อความจะใช้โครงข่ายประสาทเทียมแบบวนซ้ำหรือที่เรียกว่า Recurrent Neural Network (RNN) เนื่องจากเป็นโครงข่ายที่มีหน้าที่ในการประมวลผลข้อมูลตามลำดับทำให้สามารถจดจำการคำนวณข้อมูลก่อนหน้า ปัจจุบัน RNN ได้รับความนิยมจึงมีการพัฒนาต่อยอดทำให้เกิดโมเดลใหม่ๆ เช่น โครงข่าย Long-Short Term Memory (LSTM) หรือจะเป็น GRU เป็นต้นนอกจากนี้ยังมีงานวิจัยเกี่ยวกับการใช้งานทำนายข้อความจากความรู้สึกซึ่งมักจะใช้โครงข่าย LSTM เพื่อคัดแยกข้อความซึ่งประกอบด้วยอินพุต (ชุดข้อมูลที่เป็นข้อความ) ชั้นของ Embedding Matrix ซึ่งเป็นส่วนชั้นที่ต้องรับข้อมูลอินพุตผ่านการเตรียมข้อมูลแล้ว โดยชั้นเลเยอร์นี้จะต้องกำหนดให้มีตัวกรองเพื่อคัดคุณลักษณะตามหลักการทำงานของโครงข่ายประสาทเทียมแบบ Convolutional Neural Network (CNN) เมื่อผ่านกระบวนการนี้แล้วขั้นถัดไปจะเป็นชั้น LSTM และส่วนสุดท้ายจะเป็นโครงข่ายแบบเชื่อมต่อกันหมด (Fully connected network) ที่มีชั้นสุดท้ายเป็นคลาสที่ต้องทำนายกำหนดให้เป็นความรู้สึกด้านบวก และความรู้สึกลบ

III. ขั้นตอนวิธีที่นำเสนอ

โดยทั่วไป การวิเคราะห์ความรู้สึกของความคิดเห็นเชิงบวกและเชิงลบเป็นกระบวนการดึงข้อมูลเกี่ยวกับการระบุตัวตนแบบอัตโนมัติจากข้อความที่สร้างขึ้นโดยผู้ใช้ ซึ่งอาจจะเป็นการจำแนกในระดับประโยคหรือในระดับเอกสาร และในปัจจุบันอาจจะใช้การวิเคราะห์ด้วยการเรียนรู้เชิงลึกโดยมีกระบวนการเพื่อเตรียมชุดข้อมูลซึ่งอาจจะใช้งานอัลกอริทึมที่มีอยู่ในปัจจุบันเช่น TF-IDF, CBOW, Embedding, Skip-gram ฯลฯ ดังที่กล่าวไว้ข้างต้น โดยในงานวิจัยนี้ได้นำเทคนิคของ Word embedding ซึ่งเป็นการดำเนินการทางด้านการประมวลผลภาษา และเพื่อเพิ่มประสิทธิภาพและความเร็วในการประมวลผลต่อผู้ใช้งาน งานวิจัยนี้จึงนำเสนอระบบการวิเคราะห์ความรู้สึกโดยใช้ข้อมูลจากความคิดเห็นของผู้ใช้งานบนโซเชียลเน็ตเวิร์กโดยการนำการประมวลผลภาษาธรรมชาติ (NLP) ผสมกับอัลกอริทึมของโมเดลการเรียนรู้เชิงลึกเพื่อแก้ปัญหาการวิเคราะห์ความรู้สึกสำหรับการคัดแยกความรู้สึกเชิงบวก และเชิงลบ อีกทั้งยังเพิ่มส่วนของการนำโมเดลที่ได้รับการฝึกสอนแล้วไปใช้ทำนายผลดังรูปที่ IV



รูปที่ IV. แผนภาพแสดงการดำเนินงานของระบบ

IV. การทดลอง

ในส่วนนี้จะกล่าวถึงการทดสอบของระบบการคัดแยกความรู้สึกเชิงบวกและลบของข้อความคำวิจารณ์โดยมีขั้นตอนของการดำเนินงานดังรายละเอียดของแต่ละขั้นตอนต่อไปนี้

IV.1 การเตรียมชุดข้อมูลซึ่งประกอบด้วยชุดข้อมูลจากแหล่งต่างๆ ได้แก่

- Sentiment140 จาก Stanford University [12] เป็นข้อความทวีตเกี่ยวกับผลิตภัณฑ์จำนวน 1.5 ล้านกว่าทวีตและได้ระบุความรู้สึกของผู้วิจารณ์ด้านลบ และด้านบวก
- Tweets Airline [13] เป็นชุดข้อมูลทวีตที่ความคิดเห็นของผู้ใช้เกี่ยวกับสายการบินของสหรัฐฯ มีการรวบรวมข้อมูลในเดือนกุมภาพันธ์ 2015 มีตัวอย่าง 1 หมื่นกว่าตัวอย่าง และแบ่งออกเป็นความรู้สึกเชิงลบ เป็นกลาง และเชิงบวก ซึ่งในการทดสอบนี้ได้ตัดแบ่งนำมาใช้เฉพาะส่วนของข้อมูลการวิจารณ์ที่มีความรู้สึกเป็นบวกและลบเท่านั้น
- บทวิจารณ์ภาพยนตร์ IMDB ได้มาจาก Stanford University [14] ชุดข้อมูลนี้มีความคิดเห็นจากผู้ชมเกี่ยวกับเรื่องราวของภาพยนตร์ มี 50,000 กว่าตัวอย่าง ซึ่งแบ่งเป็นบวกและลบเช่นกัน

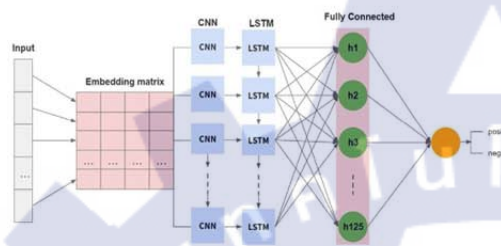
โดยทั้ง 3 ชุดข้อมูลได้เตรียมและจัดแบ่งบางส่วน ซึ่งสามารถแสดงตัวอย่างบางส่วนข้อความวิจารณ์ดังรูปที่ V

id	sentiment	review
5814_8	1	With all this stuff going down at the moment w...
2381_9	1	\The Classic War of the Worlds\ by Timothy Hi...
7759_3	0	The film starts with a manager (Nicholas Bell)...
3630_4	0	It must be assumed that those who praised this...
9495_8	1	Superbly trashy and wondrously unpretentious 8...
...	...	...
3471_3	0	this movie had more holes than a piece of swis...
6024_10	1	Last November I had a chance to see this fil...

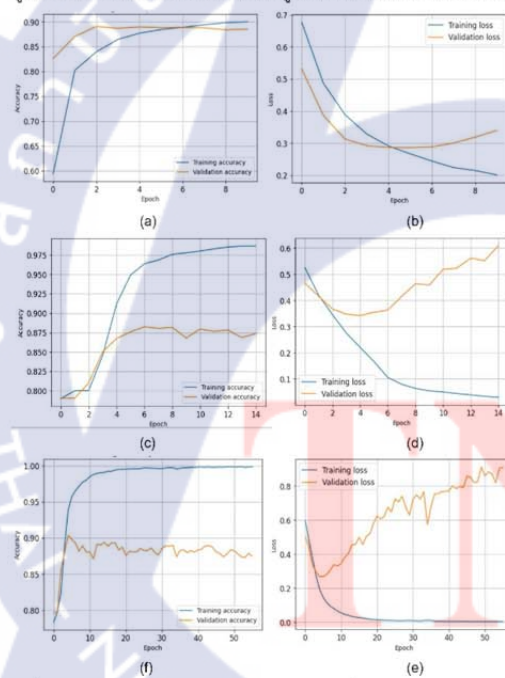
รูปที่ V. ตัวอย่างข้อมูลดิบสำหรับการทดลอง

IV.II โมเดลของการเรียนรู้เชิงลึก

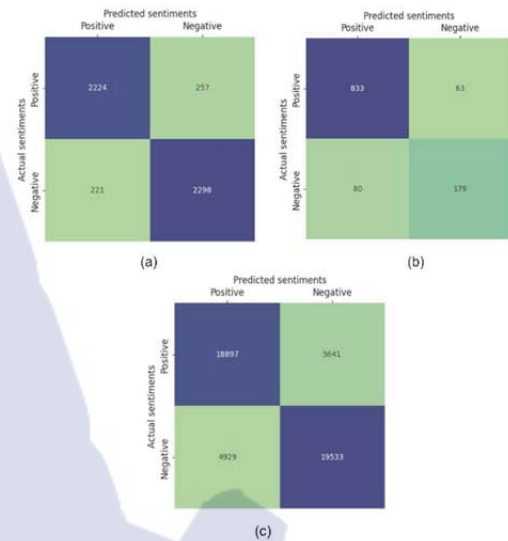
ขั้นตอนวิธี (อัลกอริทึม) ของการพยากรณ์ความรู้สึกจากข้อความวิจารณ์ที่นำเสนอมีรายละเอียดดังต่อไปนี้ ประการแรกคือ ขั้นตอนของกระบวนการเตรียมข้อมูลซึ่งประกอบด้วยการทำงานสะอาดข้อมูล (การแยกคำ การแยกประโยค การกรองตัวอักษรและลบอักขระต่างๆ และการปรับค่าให้อยู่ในรูปรากของคำศัพท์นั้นๆ ตลอดจนการกำจัดคำศัพท์ที่ปรากฏน้อยออกไป) โดยหลังจากเตรียมชุดข้อมูลเป็นที่เรียบร้อยแล้ว ประการต่อมาคือ การแปลงข้อความวิจารณ์เป็นเวกเตอร์ (หรือที่เรียกว่า Feature Vector) โดยใช้วิธี Word Embedding ซึ่งทำให้ได้ผลลัพธ์เวกเตอร์คุณลักษณะที่พร้อมจะนำไปป้อนให้กับอัลกอริทึมการเรียนรู้เชิงลึกที่ได้นำเสนอ ซึ่งในส่วนของโครงสร้างโมเดลที่นำเสนอแสดงดังรูปที่ VI



รูปที่ VI. โครงสร้างของโมเดลการเรียนรู้เชิงลึกแบบผสมผสานที่นำเสนอ



รูปที่ VII กราฟของการทดลองของอัลกอริทึมที่นำเสนอด้วย (a,b) ชุดข้อมูล Sentiment140 (c,d) ชุดข้อมูล Tweets Airline และ (e,f) ชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB



รูปที่ VIII ตาราง Confusion Matrix (a) ชุดข้อมูล Sentiment140 (b) ชุดข้อมูล Tweets Airline และ (c) ชุดข้อมูล บทวิจารณ์ภาพยนตร์ IMDB

IV.III ผลลัพธ์จากการทดลอง หลังจากที่ได้ออกแบบอัลกอริทึมที่นำเสนอสำหรับการคัดแยกอารมณ์ความรู้สึกจากข้อความวิจารณ์ดังกล่าวไว้ในหัวข้อที่ผ่านมาแล้วนั้น ในขั้นตอนนี้เป็นการดำเนินการฝึกสอนและทดสอบเพื่อประเมินประสิทธิภาพ การทดลองนี้ได้กำหนดค่าพารามิเตอร์ต่างๆ ที่สำคัญสำหรับการฝึกสอนได้แก่ 1) รอบของการฝึกสอน (epoch) เป็น 10 15 และ 50 รอบตามลำดับ 2) ชุดแบ่งข้อมูล (batch size) เท่ากับ 256, ค่าตัวแปร loss = 'binary_crossentropy', และ 3) สุดท้ายตัวแปรของขั้นตอนการเรียนรู้ (optimizer) กำหนดให้เป็น 'adam' หลังจากที่ได้ดำเนินการฝึกสอนเป็นที่เรียบร้อยแล้วผลลัพธ์สามารถแสดงกราฟของการลดลงของค่าผิดพลาด (Loss) และค่าความถูกต้อง (Accuracy)

ตารางที่ I. ผลการทำนายการคัดแยกอารมณ์ความรู้สึกบวกและลบ

	ร้อยละความถูกต้องของการทดลองด้วยอัลกอริทึมการเรียนรู้เชิงลึกแบบต่างๆ				
	RNN	LSTM	GRU	CNN-LSTM	Proposed Model
Sentiment140[12]	78.17%	78.37%	78.39%	72.58%	89.45%
Tweets Airline [14]	91.04%	89.78%	90.39%	88.66%	98.41%
IMDB review [13]	84.00%	88.08%	86.36%	88.42%	99.67%

ในการดำเนินการทดลองทางผู้วิจัยได้ออกแบบอัลกอริทึมที่นำเสนอฝึกสอนและทดสอบด้วยชุดข้อมูลที่ได้เตรียมมาตั้งแต่กล่าวในหัวข้อ IV.I ซึ่งได้แก่ชุดข้อมูล Sentiment140 ชุดข้อมูล Tweets Airline และชุดข้อมูลบทวิจารณ์ภาพยนตร์ IMDB โดยแต่ละการทดลองนั้นได้ดำเนินการใช้ชุดข้อมูลแต่ละแหล่งข้างต้นแยกออกจากกัน ซึ่งสามารถแสดงผลการฝึกสอนได้ดังรูปที่ VII(a,b) รูปที่ VII(c,d) และรูปที่ VII(e,f) สำหรับการฝึกสอนของชุดข้อมูล Sentiment140 ชุดข้อมูลบทวิจารณ์ภาพยนตร์ IMDB และ ชุดข้อมูล Tweets Airline ตามลำดับ

จากกราฟผลการฝึกสอนสามารถตั้งข้อสังเกตว่า จำนวนรอบของการฝึกสอนไม่ควรเกิน 10 รอบเนื่องจากเมื่อทดสอบเกิน 10 รอบมีแนวโน้มค่าของความถูกต้องและค่าผิดพลาดของข้อมูลทดสอบ (Validation data) นั้นล้นออก ดังแสดงในรูปที่ VII(e) และเพื่อให้เห็นถึงประสิทธิภาพของอัลกอริทึมทางผู้วิจัยได้แสดงผลลัพธ์ของการทดลองที่อยู่ในรูปของ Confusion Matrix เพื่อป้องกันข้อผิดพลาดสำหรับการตีความการตีความความรู้สึกจากข้อความที่น่าเสนอนี้ มีความสามารถในการฝึกสอนให้สามารถตัดแยกความรู้สึกจากข้อความได้เป็นอย่างดีดังรูปที่ VIII(a) รูปที่ VIII(b) และรูปที่ VIII(c) สำหรับการฝึกสอนของชุดข้อมูล Sentiment140 ชุดข้อมูลบทวิจารณ์ภาพยนตร์ IMDB และ ชุดข้อมูล Tweets Airline ตามลำดับ และในการทดลองนี้ยังได้เปรียบเทียบประสิทธิภาพของโมเดลกับโมเดลอื่นๆ ดังตารางที่ 1 ซึ่งจากผลที่แสดงในตารางนั้นจะเห็นว่าอัลกอริทึมที่น่าเสนอนี้มีความสามารถในการฝึกสอนให้สามารถตัดแยกความรู้สึกจากข้อความเฉลี่ยร้อยละ 95.84 จากตัวอย่างของการทดสอบทั้งหมด จากนั้นนำโมเดลฝึกสอน (Trained model) เป็นเรียบร้อยแล้วมาใช้ในการทำนายผลจากการนำข้อมูลจากคำวิจารณ์ในโลกโซเชียลจริง ในที่นี้ได้นำคำวิจารณ์จาก YouTube ซึ่งผลลัพธ์ของการทำนายความถูกต้องค่อนข้างดีคืออยู่ที่ร้อยละ 70

V. สรุป

งานวิจัยนี้ตระหนักเห็นถึงความจำเป็นอย่างยิ่งที่จะต้องพัฒนางานวิจัยในด้านการประมวลผลภาษาธรรมชาติ (NLP) โดยนำโมเดลการเรียนรู้เชิงลึกมาใช้เพื่อแก้ปัญหาการตัดแยกความรู้สึกเชิงบวก และเชิงลบจากชุดข้อมูลที่เป็นข้อความวิจารณ์จากผู้บริโภคซึ่งอาจนำไปใช้ในการปรับปรุงผลิตภัณฑ์หรือสินค้าต่างๆ โดยในงานวิจัยนี้ได้ออกแบบขั้นตอนวิธี (อัลกอริทึม) ของการพยากรณ์ความรู้สึกจากข้อความวิจารณ์ที่ประกอบด้วยส่วนสำคัญอยู่ 2 ส่วนดังนี้ ซึ่งประกอบด้วยการทำความเข้าใจข้อมูล (การแยกคำ การแยกประโยค การกรองตัวอักษรและลบอักขระต่างๆ และการปรับค่าให้อยู่ในรูปรากของค่าศัพท์นั้นๆ ตลอดจนการกำจัดคำศัพท์ที่ปรากฏน้อยออกไป) โดยหลังจากเตรียมชุดข้อมูลเป็นเรียบร้อยแล้วแล้วประการต่อมาคือการแปลงข้อความวิจารณ์เป็นเวกเตอร์ (หรือที่เรียกว่า Feature Vector) โดยใช้วิธี Word Embedding ซึ่งทำให้ได้ผลลัพธ์เวกเตอร์คุณลักษณะที่พร้อมจะนำไปป้อนให้กับอัลกอริทึมการเรียนรู้เชิงลึกที่ได้นำเสนอ โดยผู้วิจัยได้ดำเนินการทดลองและทดสอบด้วยอัลกอริทึมที่น่าเสนอเพื่อฝึกสอนและทดสอบด้วยชุดข้อมูล 3 ชุดได้แก่ ชุดข้อมูล Sentiment140 ชุดข้อมูล Tweets Airline และชุดข้อมูลบทวิจารณ์ภาพยนตร์ IMDB โดยแต่ละการทดลองนั้นได้ดำเนินการใช้ชุดข้อมูลแต่ละแหล่งข้างต้นแยกออกจากกัน และจากผลการทดลองจะได้ว่าอัลกอริทึมที่น่าเสนอนี้มีความสามารถในการฝึกสอนให้สามารถตัดแยกความรู้สึกจากข้อความเฉลี่ยในระดับที่ดีมากจากตัวอย่างของการทดสอบทั้งหมด

สำหรับการพัฒนาในอนาคตผู้วิจัยคาดว่าจะพัฒนาอัลกอริทึมให้มีความสามารถทำนายอารมณ์ความรู้สึกจากคำวิจารณ์โซเชียลให้ได้ดีมีประสิทธิภาพมากขึ้น และนำไปใช้งานได้หลากหลายและกว้างขวางมากขึ้น

เอกสารอ้างอิง

- [1] M. Lukauskas, V. Pilinkienė, J. Bruneckienė, A. Stundzienė, A. Grybauskas, and T. Ruzgas, "Economic Activity Forecasting Based on

the Sentiment Analysis of News," *Mathematics*, vol. 10, no. 19, p. 3461, Sep. 2022.

- [2] Viswanathan, A., Umamaheswari, M. and Mohanasundaram, R. (2022) "Sentiment analysis and accuracy prediction of medical dataset using deep adversarial networks," *Quing Publications*, pp. 41–53.
- [3] Mohanasundaram, R. (2022) "[3] Deep associative learning approach for bio-medical sentiment," *Quing Publications*, pp. 41–53.
- [4] K. Mishev, A. Gjorgjevikj, I. Vodenska, L. T. Chitkushev and D. Trajanov, "Evaluation of Sentiment Analysis in Finance: From Lexicons to Transformers," in *IEEE Access*, vol. 8, pp. 131662-131682, 2020.
- [5] Kauffmann, Peral, Gil, Ferrández, Sellers, and Mora, "Managing Marketing Decision-Making with Sentiment Analysis: An Evaluation of the Main Product Features Using Text Data Mining," *Sustainability*, vol. 11, no. 15, p. 4235, Aug. 2019.
- [6] C. -C. Lee, Z. Gao and C. -L. Tsai, "BERT-Based Stock Market Sentiment Analysis," 2020 IEEE International Conference on Consumer Electronics - Taiwan (ICCE-Taiwan), Taoyuan, Taiwan, 2020, pp. 1-2.
- [7] Grissette H, Nfaoui EH. Deep associative learning approach for bio-medical sentiment analysis utilizing unsupervised representation from large-scale patients' narratives. *Pers Ubiquitous Comput.* 2021 Aug 11:1-15.
- [8] M. Hadwan, M. Al-Sarem, F. Saeed, and M. A. Al-Hagery, "An Improved Sentiment Classification Approach for Measuring User Satisfaction toward Governmental Services' Mobile Apps Using Machine Learning Methods with Feature Engineering and SMOTE Technique," *Applied Sciences*, vol. 12, no. 11, p. 5547, May 2022.
- [9] Mohammed Rameez Qureshi, Sidharth Ranjan, Rajakrishnan Rajkumar, and Kushal Shah. 2019. A Simple Approach to Classify Fictional and Non-Fictional Genres. In *Proceedings of the Second Workshop on Storytelling*, pages 81–89, Florence, Italy. Association for Computational Linguistics.
- [10] M. U. Salur and I. Aydin, "A Novel Hybrid Deep Learning Model for Sentiment Classification," in *IEEE Access*, vol. 8, pp. 58080-58093, 2020.
- [11] Cach N. Dang, Maria N. Moreno-García, Fernando De la Prieta, "Hybrid Deep Learning Models for Sentiment Analysis", *Complexity*, vol. 2021, Article ID 9986920, 16 pages, 2021.
- [12] Available online: <http://help.sentiment140.com/site-functionality> (accessed on 12 Feb 2023).
- [13] Available online: <https://www.kaggle.com/crowdfunder/twitter-airline-sentiment> (accessed on 12 Feb 2023).
- [14] Available online: <https://www.kaggle.com/datasets/lakshmi25npathi/imdb-dataset-of-50k-movie-reviews> (accessed on 12 Feb 2023).



_____ มอบเกียรติบัตรฉบับนี้ไว้เพื่อแสดงว่า _____

คุณินา พงศ์ศักดิ์ศรี อัดนา เซนโตะ

ได้นำเสนอบทความเรื่อง

การวิเคราะห์ความคิดเห็นจากข้อความความคิดเห็นด้วยการเรียนรู้เชิงลึก

ในงานประชุมสหวิทยาการระดับชาติสถาบันเทคโนโลยีไทย-ญี่ปุ่น ครั้งที่ 9 (TNIAC2023)
จัดโดย สถาบันเทคโนโลยีไทย-ญี่ปุ่น (TNI), สมาคมส่งเสริมเทคโนโลยี (ไทย-ญี่ปุ่น) (TPA)
และสมาคมปัญญาประดิษฐ์ประเทศไทย (AIAT)
ระหว่างวันที่ 18-19 พฤษภาคม 2566



รองศาสตราจารย์ ดร. รัตติกอร์ วรากุลศิริพันธุ์
ประธานคณะกรรมการจัดประชุมวิชาการ TNIAC 2023