



การวิเคราะห์ประสบการณ์จากการใช้บริการโรงพยาบาลในประเทศไทย
จากความคิดเห็นของผู้ใช้บริการ

CUSTOMER EXPERIENCE ANALYTICS FOR THAILAND HOSPITALS



พนิดา กาศกลางดอน

บัณฑิตวิทยาลัย มหาวิทยาลัยศรีนครินทรวิโรฒ

2563

การวิเคราะห์ประสบการณ์จากการใช้บริการโรงพยาบาลในประเทศไทย
จากความคิดเห็นของผู้ใช้บริการ



พนิดา กาศกลางดอน

สารนิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตร
วิทยาศาสตรมหาบัณฑิต สาขาวิชาวิทยาการข้อมูล
คณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ
ปีการศึกษา 2563
ลิขสิทธิ์ของมหาวิทยาลัยศรีนครินทรวิโรฒ

CUSTOMER EXPERIENCE ANALYTICS FOR THAILAND HOSPITALS



PANIDA KATKLANGDON

A Master's Project Submitted in Partial Fulfillment of the Requirements

for the Degree of MASTER OF SCIENCE

(Data Science)

Faculty of Science, Srinakharinwirot University

2020

Copyright of Srinakharinwirot University

สารนิพนธ์

เรื่อง

การวิเคราะห์ประสบการณ์จากการใช้บริการโรงพยาบาลในประเทศไทย
จากความคิดเห็นของผู้ใช้บริการ

ของ

พนิดา กาศกลางดอน

ได้รับอนุมัติจากบัณฑิตวิทยาลัยให้นับเป็นส่วนหนึ่งของการศึกษาตามหลักสูตร

ปริญญาวิทยาศาสตรมหาบัณฑิต สาขาวิชาวิทยาการข้อมูล

ของมหาวิทยาลัยศรีนครินทรวิโรฒ

(รองศาสตราจารย์ นายแพทย์ฉัตรชัย เอกปัญญาสกุล)

คณบดีบัณฑิตวิทยาลัย

คณะกรรมการสอบปากเปล่าสารนิพนธ์

..... ที่ปรึกษาหลัก

(อาจารย์ ดร.ศิริสรรพ เหล่าหะเกียรติ)

..... ประธาน

(ผู้ช่วยศาสตราจารย์ ดร.อัศรินทร์ ไพบูลย์พานิช)

..... ที่ปรึกษาร่วม

(อาจารย์ ดร.รัตน์ชัยนันท์ ธรรมสุจริต)

..... กรรมการ

(ผู้ช่วยศาสตราจารย์ ดร.นุรีย์ วิวัฒน์วัฒนา)

ชื่อเรื่อง การวิเคราะห์ประสบการณ์จากการใช้บริการโรงพยาบาลในประเทศไทย
จากความคิดเห็นของผู้ใช้บริการ

ผู้วิจัย พนิดา กาศกลางดอน
ปริญญา วิทยาศาสตร์มหาบัณฑิต
ปีการศึกษา 2563
อาจารย์ที่ปรึกษา อาจารย์ ดร. ศิริสรรพ เหล่าหะเกียรติ
อาจารย์ที่ปรึกษาร่วม อาจารย์ ดร. รัตน์ชัยนันท์ ธรรมสุจริต

การรับรู้ถึงปัญหาและการแก้ปัญหาที่ไม่ตรงจุดในการตอบสนองความต้องการของผู้ใช้บริการได้
ทันท่วงที ถือเป็นปัญหาสำคัญที่ส่งผลกระทบต่อคุณภาพการให้บริการของโรงพยาบาลในปัจจุบัน สาเหตุหลักของ
ปัญหา เนื่องจากมีผู้ให้บริการจำนวนมากที่ได้รับปัญหาจากการบริการในหลายๆด้านของโรงพยาบาล และด้วย
วิธีการต่างๆ ในการรวบรวมข้อมูลของผู้ใช้บริการและการได้มาซึ่งข้อมูลล่าช้าซึ่งไม่ได้สะท้อนถึงปัญหาที่แท้จริง
งานวิจัยนี้นำเสนอแนวทางแก้ไขปัญหาดังกล่าวโดยรวบรวมความคิดเห็นของผู้มีประสบการณ์การให้บริการจากเว็บ
บอร์ด ซึ่งรวดเร็วและได้รับข้อมูลที่แท้จริงจากการรีวิว ข้อมูลเหล่านี้จะถูกวิเคราะห์เพื่อหาความรู้เชิงลึกโดยการใช้
โมเดลการเรียนรู้ของเครื่อง ด้วยข้อมูลข้อความจากประสบการณ์ผู้ป่วย และใช้ข้อมูลที่รวบรวมมาสร้างรายงาน
ผลผ่านแดชบอร์ด ผลการเปรียบเทียบใช้การทดสอบการจำลองโดยใช้เทคนิคอัลกอริทึม 3 เทคนิค คือ นาร์วีย์ฟ
เบย์ (Naïve Bayes), การสุ่มป่าไม้ (Random Forest) และ โครงข่ายประสาทเทียมแบบหน่วยความจำระยะสั้น
แบบยาว (Long-Short-Term Memory) จากผลการทดสอบพบว่าอัลกอริทึม Random Forest มีประสิทธิภาพ
สูงสุดในการทำนายความคิดเห็นเชิงลบโดยมีค่าการเรียกคืน (Recall) ที่ 0.71 และผลวิเคราะห์จากการนำข้อมูล
ของโรงพยาบาลรัฐและเอกชนเพื่อหาข้อมูลเชิงลึก ในภาพรวมการวิเคราะห์ของระบบจัดการคุณภาพ
โรงพยาบาลบนรายงาน พบว่าในด้านบุคลากร ด้านการจัดการโครงสร้าง ด้านการดำเนินงาน ด้านการเงิน และ
ด้านอื่นๆ ของโรงพยาบาลรัฐถูกพูดถึงในเชิงลบมากกว่าโรงพยาบาลเอกชนทั้ง 5 ด้าน ด้านที่ชัดเจนที่สุดคือ
โครงสร้างโรงพยาบาล

คำสำคัญ : การวิเคราะห์ข้อความ, การรีวิวประสบการณ์การให้บริการโรงพยาบาล, การเรียนรู้ของเครื่อง, ข้อมูล
เชิงลึก, การจำแนกประเภทความรู้สึก, เทคนิคการสุ่มป่าไม้

Title	CUSTOMER EXPERIENCE ANALYTICS FOR THAILAND HOSPITALS
Author	PANIDA KATKLANGDON
Degree	MASTER OF SCIENCE
Academic Year	2020
Thesis Advisor	Professor Dr. Sirisup Laohakiat
Co Advisor	Professor Dr. Ratchainant Thammasudjarit

Recognizing and solving problems that do not meet the needs of users in a timely manner is an important problem that affects the quality of hospital services at present. This problem was caused due to a large number of users who have had problems with services in many areas of the hospital and in various ways for the collection of user information and a delay in obtaining information that does not reflect the real problem. This research presents a solution to the problem by gathering reviews of those who have experienced using the service from the web board, which is fast and gets real information from reviews. These data analyzed insight information with machine learning models with text data from reviews of patient experience and use the collected and analyzed data to generate reports through the dashboard. The comparison results were based on simulation testing using three algorithmic techniques: Naïve Baye, Random Forest, Long and Short Term Memory (LSTM) found that the Random Forest algorithm was the most effective at predicting negative reviews with recall at 0.71 and the results of analyzing the use of data from public and private hospitals to find insight information In the overview of the hospital quality management system through reports, found that in staff, service, infrastructure, processing, finance and other public hospitals were discussed more negatively than private hospitals in all six areas. The most obvious aspect was the hospital structure.

Keyword : Text Analytics, Customer Experience of Hospitals, Machine Learning, Insight Data, Sentiment Classification, Random Forest

กิตติกรรมประกาศ

สารนิพนธ์ฉบับนี้เสร็จสมบูรณ์เป็นอย่างดีได้ด้วยความช่วยเหลือ และการให้คำปรึกษาจาก อาจารย์ที่ปรึกษา ได้แก่ อ.ดร. รัตน์ชัยนันท์ ธรรมสุจริต อ.ดร. โสภณ มงคลลักษณ์ อ.ดร. ศิริสรวพ เหล่าหะเกียรติ และ คณะอาจารย์ที่เป็นคณะกรรมการสอบสารนิพนธ์ทุกท่าน สำหรับคำแนะนำในทุกขั้นตอนที่ได้ทำการศึกษา รายวิชาสารนิพนธ์ ตั้งแต่การวางแผนการดำเนินงาน จนนำเสนอ ผลงานวิจัยที่สำเร็จลุล่วงไปได้ด้วยดี ตลอดจนการเขียนรูปเล่มรายงาน การตรวจสอบ และการแก้ไข ข้อบกพร่องต่างๆในงานวิจัยนี้ฉบับนี้จนเสร็จสมบูรณ์ ขอขอบคุณคณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ อันเป็นสถานศึกษาที่ประสิทธิ์ ประสาทวิชาความรู้ขอกราบขอบพระคุณคณาจารย์ ประจำภาควิทยาการคอมพิวเตอร์ คณะวิทยาศาสตร์ มหาวิทยาลัยศรีนครินทรวิโรฒ ทุกท่านที่ได้ให้การสั่งสอนรายวิชาที่เป็นพื้นฐานในการศึกษาสาขาวิทยาการข้อมูล ให้คำแนะนำและข้อคิดเห็นอันเป็นประโยชน์ต่อการทำสารนิพนธ์ขอขอบคุณ เจ้าหน้าที่และบุคลากรคณะวิทยาศาสตร์ทุกท่าน ที่ได้ให้ความช่วยเหลือและอำนวยความสะดวก สะดวกในด้านต่างๆ ขอขอบคุณเพื่อนนักศึกษาคณะ วิทยาศาสตร์ สาขาวิทยาการข้อมูล ทุกท่านที่ให้คำแนะนำ ช่วยเหลือและเป็นกำลังใจในการศึกษา ตลอดมา สุดท้ายผลอันจะเป็นประโยชน์ความดีความงามทั้งปวง ที่เกิดขึ้นจากการศึกษาสารนิพนธ์นี้ ขอมอบแด่คุณพ่อและคุณแม่ที่เคารพยิ่งและหากมีข้อบกพร่องด้วยประการใดๆ ผู้วิจัยขอน้อมรับไว้ ด้วยความขอบคุณยิ่ง

พนิดา กาศกลางดอน

สารบัญ

	หน้า
บทคัดย่อภาษาไทย.....	ง
บทคัดย่อภาษาอังกฤษ.....	จ
กิตติกรรมประกาศ.....	ฉ
สารบัญ.....	ช
สารบัญตาราง.....	ญ
สารบัญรูปภาพ.....	ฎ
บทที่ 1 บทนำ.....	1
1.1 ความสำคัญและความเป็นมาของงานวิจัย.....	1
1.2 วัตถุประสงค์ของการงานวิจัย.....	3
1.3 ขอบเขตของการวิจัย.....	3
1.4 ประโยชน์ที่คาดว่าจะได้รับจากงานวิจัย.....	4
บทที่ 2 เอกสารทฤษฎีและงานวิจัยที่เกี่ยวข้อง.....	5
2.1 งานวิจัยที่เกี่ยวข้อง.....	5
2.2 เอกสารทฤษฎีที่เกี่ยวข้อง.....	7
2.2.1 การวิเคราะห์ความรู้สึก.....	7
2.2.2 การเก็บรวบรวมข้อมูลจากเว็บไซต์.....	7
2.2.3 การสกัดข้อมูลจากเว็บไซต์.....	8
2.2.4 การประมวลผลภาษาธรรมชาติ.....	8
2.2.5 การสร้างคุณลักษณะของข้อมูล.....	9
2.2.5.1 การสร้างคุณลักษณะด้วยเทคนิค TFIDF.....	10
2.2.5.2 การสร้างคุณลักษณะด้วยเทคนิค Word Embedding.....	11

2.2.6 เทคนิคการเรียนรู้ของเครื่อง	12
2.2.7 อัลกอริทึม Multinomial Naïve Bayes	12
2.2.8 อัลกอริทึม Long-short-term memory	13
2.2.9 อัลกอริทึม Random Forest	13
2.2.10 เทคนิคการแสดงผลข้อมูลของคำ	15
2.2.11 แนวคิดในการออกแบบแดชบอร์ด	15
บทที่ 3 วิธีการดำเนินงานวิจัย	17
3.1 การวางแผนการดำเนินงานวิจัย	17
3.2 การเก็บรวบรวมข้อมูล	18
3.3 การประมวลผลภาษาธรรมชาติ	21
3.4 การกำหนดประเภทความรู้สึก	22
3.5 การสร้างคุณลักษณะเพื่อการแปลงข้อมูล	23
3.6 การแบ่งข้อมูลสำหรับชุดฝึกฝนและชุดทดสอบ	23
3.7 การสร้างแบบจำลองการแยกประเภทความรู้สึก	24
3.8 การประเมินผลประสิทธิภาพแบบจำลองการเรียนรู้ของเครื่อง	29
3.9 การประยุกต์ใช้แบบจำลองการเรียนรู้เพื่อแยกประเภทความรู้สึกและหาข้อมูลเชิงลึก	30
3.9.1 ข้อมูลรีวิวกจากประสบการณ์การใช้บริการของโรงพยาบาลรัฐและเอกชน	31
3.9.2 การแยกประเภทความรู้สึกด้วยแบบจำลองการเรียนรู้ของเครื่อง	31
3.9.3 การจำแนกหัวข้อตัวชี้วัดของข้อมูลรีวิวก	31
3.9.4 การแสดงผลรายงานผ่านแดชบอร์ด	32
บทที่ 4 ผลการดำเนินการวิจัย	35
4.1 การวิเคราะห์จำนวนข้อมูลรีวิวกเชิงลบและเชิงบวกบนกราฟแท่ง	38
4.2 การวิเคราะห์ข้อมูลสัดส่วนของหัวข้อตัวชี้วัดบนกราฟพาย	39

4.3 การวิเคราะห์จำนวนข้อมูลรีวิวเชิงบวกและเชิงลบในช่วงเวลารายเดือนบนกราฟเส้น.....	40
4.4. การวิเคราะห์คำที่ปรากฏในข้อความรีวิวแสดงผ่านเวิร์ดคลาวด์.....	41
4.6 การวิเคราะห์จำนวนข้อมูลรีวิวเชิงลบและเชิงบวกในแต่ละหัวข้อตัวชี้วัด	42
บทที่ 5 สรุปผลการวิจัย อภิปรายผล และข้อเสนอแนะ	44
5.1 สรุปผลการวิจัย	44
5.2 การอภิปรายผล	45
5.3 ข้อจำกัดและอุปสรรคของงานวิจัย	47
5.4 ข้อเสนอแนะ.....	47
บรรณานุกรม.....	49
ภาคผนวก.....	50
ภาคผนวก ก.....	51
ภาคผนวก ข.....	54
ภาคผนวก ค	57
ภาคผนวก ง	59
ประวัติผู้เขียน.....	61

สารบัญตาราง

	หน้า
ตาราง 1 ตารางตัวอย่างการคำนวณความถี่ของคำที่ปรากฏ(TF Calculation).....	9
ตาราง 2 ตารางตัวอย่างการคำนวณค่าความถี่ผกผัน(IDF Calculation)	9
ตาราง 3 ตารางตัวอย่างการคำนวณTFIDF	10
ตาราง 4 รายละเอียดข้อมูลที่ได้รับจากการสกัดข้อมูล	20
ตาราง 5 ตัวอย่างการแปลภาษาของข้อมูลรีวิว	21
ตาราง 6 ตารางคลังข้อมูลรีวิวที่จัดประเภทความรู้สึก	23
ตาราง 7 คลังข้อมูลที่แบ่งชุดข้อมูลฝึกฝนและชุดข้อมูลทดสอบ	24
ตาราง 8 ข้อมูลแสดงคุณภาพของข้อมูลจากแหล่งข้อมูล.....	24
ตาราง 9 ตารางแสดงค่าพารามิเตอร์ของการสร้างแบบจำลองด้วยอัลกอริทึม Random Forest ..	26
ตาราง 10 ตารางแสดงค่าพารามิเตอร์ของการสร้างแบบจำลองด้วยอัลกอริทึม Naive Bayes ...	27
ตาราง 11 ตารางแสดงค่าพารามิเตอร์ของการสร้างแบบจำลองด้วยอัลกอริทึม LSTM.....	28
ตาราง 12 ตารางผลลัพธ์ Summary ของโมเดล LSTM ที่ได้จากการใช้พารามิเตอร์ข้างต้น.....	28
ตาราง 13 ข้อมูลรีวิวสำหรับการวิเคราะห์ความรู้สึกจากประสบการณ์การใช้บริการโรงพยาบาลในประเทศไทย	31
ตาราง 14 ตัวอย่างกลุ่มคำที่เกี่ยวข้องในหัวข้อตัวชี้วัด.....	32
ตาราง 15 ตารางรายงานผลค่าทางสถิติ Accuracy, Precision, Recall และ F1-Score	36
ตาราง 16 ข้อมูลจำนวนรีวิวเชิงลบและรีวิวเชิงบวกในหัวข้อตัวชี้วัดของโรงพยาบาลรัฐและเอกชน	43

สารบัญรูปภาพ

หน้า

ภาพประกอบ 1 แนวคิดการสร้าง One Hot Encoding	11
ภาพประกอบ 2 หลักการสร้างแบบจำลองการเรียนรู้ด้วยอัลกอริทึม LSTM	13
ภาพประกอบ 3 หลักการสร้างแบบจำลองการเรียนรู้ด้วยอัลกอริทึม Random Forest	14
ภาพประกอบ 4 แผนผังการดำเนินงานวิจัย	17
ภาพประกอบ 5 หน้าเว็บไซต์แสดงรายชื่อโรงพยาบาลที่ได้รับการรีวิว	18
ภาพประกอบ 6 ตัวอย่างข้อมูลรีวิวบนเว็บไซต์	19
ภาพประกอบ 7 แผนผังการประมวลผลภาษาธรรมชาติ	21
ภาพประกอบ 8 แผนผังกระบวนการทำ Sentiment Analysis	30
ภาพประกอบ 9 แดชบอร์ดการแสดงผลรายงานการวิเคราะห์รีวิวจากประสบการณ์ของผู้ใช้บริการ	32
ภาพประกอบ 10 แดชบอร์ดการแสดงผล word cloud ในหัวข้อตัวชี้วัดแต่ละด้าน	33
ภาพประกอบ 11 ผลการแสดงผลรายงานการวิเคราะห์ข้อมูลรีวิวเชิงของโรงพยาบาลรัฐ	37
ภาพประกอบ 12 ผลการแสดงผลรายงานการวิเคราะห์ข้อมูลรีวิวเชิงลบของเอกชน	37
ภาพประกอบ 13 กราฟแท่งแสดงสัดส่วนรีวิวเชิงลบและเชิงบวกของโรงพยาบาลเอกชน	38
ภาพประกอบ 14 กราฟแท่งแสดงสัดส่วนรีวิวเชิงลบและเชิงบวกของโรงพยาบาลรัฐ	38
ภาพประกอบ 15 กราฟพายแสดงสัดส่วนรีวิวของโรงพยาบาลรัฐตามหัวข้อตัวชี้วัด	39
ภาพประกอบ 16 กราฟพายแสดงสัดส่วนรีวิวของโรงพยาบาลเอกชนตามหัวข้อตัวชี้วัด	39
ภาพประกอบ 17 กราฟเส้นแสดงจำนวนรีวิวเชิงบวกและลบของโรงพยาบาลเอกชน	40
ภาพประกอบ 18 กราฟเส้นแสดงจำนวนรีวิวเชิงบวกและลบของโรงพยาบาลรัฐบาล	40
ภาพประกอบ 19 เวิร์ดคลาวด์แสดงความถี่คำที่ปรากฏในข้อความรีวิวของโรงพยาบาลเอกชน ..	41
ภาพประกอบ 20 เวิร์ดคลาวด์แสดงความถี่คำที่ปรากฏในข้อความรีวิวของโรงพยาบาลรัฐบาล ..	41

ภาพประกอบ 21 กราฟแท่งแสดงจำนวนข้อมูลรีวิวเชิงลบและเชิงบวกของโรงพยาบาลเอกชน...42	
ภาพประกอบ 22 กราฟแท่งแสดงจำนวนข้อมูลรีวิวเชิงลบและเชิงบวกของโรงพยาบาลรัฐ.....42	



บทที่ 1

บทนำ

1.1 ความสำคัญและความเป็นมาของงานวิจัย

ในยุคดิจิทัลนี้ธุรกิจหรือผู้ประกอบการได้หันมาการใช้เทคโนโลยี เพื่อช่วยให้ระบบการบริการดีขึ้นและตอบสนองความต้องการของลูกค้าได้ทันทั่วทั้งที่โดยที่ผู้ประกอบการสามารถตัดสินใจจากข้อมูลและต่อยอดธุรกิจจากการใช้ประโยชน์จากข้อมูลของลูกค้า (Customer Data) เช่น ข้อความรีวิวการบริการหรือการซื้อสินค้า ข้อมูลพฤติกรรมการใช้บริการของลูกค้า ผ่านเว็บไซต์ หรือแอปพลิเคชัน ซึ่งข้อมูลเหล่านี้สามารถบ่งชี้ได้ว่าลูกค้ารู้สึกและมีความสนใจ กับสินค้า และการบริการอย่างไรบ้าง เมื่อผู้ประกอบการมีข้อมูลเหล่านี้จะสามารถนำมาสร้างกลยุทธ์ทางธุรกิจ เพื่อให้เกิดผลลัพธ์ที่มีประสิทธิภาพ

ปัจจุบันการแสดงความคิดเห็นเพื่อการสื่อสารและแสดงอารมณ์และความรู้สึกเกิดขึ้นมากมายในโซเชียลมีเดียบนแพลตฟอร์มต่างๆการใช้เทคโนโลยีในการวิเคราะห์ข้อมูลของลูกค้ามีมาก ในธุรกิจหลายด้านเช่น การวิเคราะห์ความรู้สึกของลูกค้าใน ธุรกิจด้านอาหาร ธุรกิจโรงแรม และธุรกิจในการแพทย์ที่มีการนำเสนอที่ค่อนข้างน้อยธุรกิจด้านต่างๆจะนำข้อมูล เหล่านี้มาวิเคราะห์ให้ถูกถ่ายทอดออกมาเป็นข้อมูลเชิงลึกที่เกี่ยวข้องกับอารมณ์และความรู้สึกของ บุคคลที่มีผลทางธุรกิจ

การให้บริการอย่างธุรกิจในการแพทย์ที่ผู้คนให้ความสนใจและมีความต้องการที่จะได้รับการบริการอย่างดีที่สุด ในปัจจุบันยังมีการวิจารณ์มากมายในเว็บไซต์ออนไลน์ ซึ่งปัญหาที่เกิดขึ้นจากการบริการยังไม่สามารถแก้ปัญหาที่เกิดขึ้นได้ เนื่องจากความล่าช้า ในการรับรู้ปัญหาที่เกิดขึ้นการวิจารณ์หรือการแสดงความคิดเห็นแบบตรงไปตรงมา (Reluctant Feedback) ที่เกิดขึ้นจากความคิดเห็นจากผู้ใช้บริการและเนื่องจากการนำข้อมูล ของผู้ใช้บริการมาใช้ประโยชน์ ซึ่งมีการดำเนินการที่ค่อนข้างล่าช้าตั้งแต่กระบวนการ จัดเก็บข้อมูลจนถึงการนำเสนอผลการประเมิน การให้บริการของระบบโรงพยาบาลต่อผู้ กำหนดนโยบายหรือผู้บริหาร จึงทำให้โรงพยาบาลตอบสนองต่อผู้ใช้บริการได้ไม่ทันทั่วทั้งที่ เพื่อแก้ปัญหาในความล่าช้าในการรับรู้ถึงปัญหาที่เกิดขึ้นจึงต้องมีการเก็บข้อมูลของผู้ใช้บริการ ให้มีความรวดเร็วและดำเนินการวิเคราะห์ให้ไวที่สุดเพื่อให้ทราบว่าผู้ใช้บริการได้รับปัญหา หรือมีปัญหาอะไรที่เกิดขึ้นส่งผลให้ผู้บริหารดำเนินการตอบสนองความต้องการของผู้ใช้บริการได้ อย่างทันทั่วทั้งที่

ความล่าช้าในการรับรู้การวิจารณ์หรือการแสดงความคิดเห็นแบบตรงไปตรงมา (Reluctant Feedback) ที่เกิดขึ้นจากความคิดเห็นจากผู้ใช้บริการ ถือเป็นปัญหาสำคัญของการ

บริหารจัดการของโรงพยาบาลในประเทศไทย ด้วยปัจจุบันระบบของโรงพยาบาลในประเทศไทยมีการเก็บข้อมูลการใช้บริการของผู้ป่วยในรูปแบบที่หลากหลาย เช่น การทำแบบสอบถามของผู้ใช้บริการโดยตรง(Paper questionnaire) การทำแบบถามออนไลน์(Online questionnaire) การรับข้อความจากกล่องรับจดหมายร้องเรียนในการเก็บข้อมูล ทำให้มีกระบวนการดำเนินการหลายขั้นตอน ซึ่งก่อให้เกิดความล่าช้าในการได้มาของข้อมูล ส่งผลให้เจ้าหน้าที่ผู้รับผิดชอบหรือผู้บริหารไม่สามารถนำข้อมูลมาวิเคราะห์เพื่อจัดการกับปัญหาและตอบสนองต่อผู้ให้บริการได้ทันเวลาที่โดยสาเหตุของความล่าช้านี้มาจากกระบวนการและลำดับการทำงานที่มีหลายขั้นตอนดังนี้

1.Data Acquisition คือการได้มาของข้อมูล ตั้งแต่อดีตจนถึงปัจจุบันการเก็บข้อมูลของโรงพยาบาลจะสามารถแบ่งออกเป็นสองรูปแบบคือ

Active Acquisition คือการส่งหน่วยงานลงไปทำการสำรวจ (survey) เพื่อให้ผู้มาใช้บริการทำแบบถามของโรงพยาบาล ซึ่งแบบสอบถามของแต่ละโรงพยาบาลจะออกแบบแตกต่างกันไป โดยที่แต่ละโรงพยาบาลมีมาตรฐานการออกแบบแบบสอบถามไม่เหมือนกัน ทำให้คุณภาพของข้อมูล (Data quality) ไม่ได้มาตรฐาน และอีกหนึ่งสาเหตุที่ทำให้ข้อมูลที่ได้ยังมีความไม่มากพอคือ ผู้ตอบแบบสอบถามอาจยังไม่มีความพร้อมในการทำแบบสอบถามเมื่อต้องกรอกแบบถามในสถานที่โรงพยาบาล

Passive Acquisition คือการรอผู้ให้บริการให้ feedback กับทางโรงพยาบาลเอง หรือถ้าไม่มี feedback จากผู้บริการก็จะไม่สามารถเก็บข้อมูลได้ ตัวอย่างเช่น การเก็บข้อมูลของผู้ใช้บริการจากกล่องส่งข้อความร้องเรียน ซึ่งจะต้องใช้เวลานานให้ข้อมูลมีปริมาณที่มากพอจึงจะสามารถเก็บข้อมูลเข้าสู่ระบบเพื่อนำไปวิเคราะห์

2.Data Entry คือการบันทึกข้อมูลลงแหล่งจัดเก็บข้อมูล เนื่องจากข้อมูลที่นำมาบันทึกลงในระบบมีหลายรูปแบบ เช่น แบบสอบถามจาก Google form แบบสอบถามทางอิเล็กทรอนิกส์ และ แบบสอบถามในรูปแบบกระดาษ ซึ่งการบันทึกข้อมูลต้องใช้เวลา และการบันทึกข้อมูลมีโอกาสที่จะตกหล่นและมีความผิดพลาดได้

3.Data Pre-processing คือการเตรียมข้อมูลให้พร้อมเพื่อนำมาวิเคราะห์ ในการทำงานปัจจุบันเมื่อหน่วยงานได้รับข้อมูลมานำมาวิเคราะห์ทางหลักทางสถิติ (Statistical Analysis) จัดหมวดหมู่ของปัญหาที่เกิดขึ้นเอง แปลความหมายของความคิดเห็นของผู้ใช้บริการเอง ซึ่งปัจจุบันยังไม่มีการใช้การวิเคราะห์แบบอัตโนมัติ เช่น การวิเคราะห์ความรู้สึก (sentiment Analysis) จึงเป็นสาเหตุที่ทำให้การประมวลผลการวิเคราะห์มีความล่าช้า

4. Virtualization Report คือการแสดงผลจากการวิเคราะห์ ข้อมูล เมื่อได้ผลการวิเคราะห์แล้ว สามารถแสดงผลผ่าน dashboard เพื่อรายการปัญหาที่เกิดขึ้นต่อผู้บริหารได้ทันที เนื่องจากวิธีการเก็บข้อมูลแบบ Common practice ของโรงพยาบาลในปัจจุบันยังไม่ถึงการสร้างระบบการจัดเก็บข้อมูลและวิเคราะห์ข้อมูลเพื่อแสดงผลแบบอัตโนมัติและทำให้ผู้บริหารทราบปัญหาได้ช้าและไม่สามารถจัดการกับปัญหาที่เกิดขึ้นได้ทันถ่วงที

ผู้วิจัยจึงสร้างแบบจำลองการทำนายความรู้สึกเพื่อวิเคราะห์ข้อความความคิดเห็นของผู้ใช้บริการโรงพยาบาลเพื่อให้ทราบถึงมุมมองของผู้ใช้บริการเมื่อเข้ามาใช้บริการที่โรงพยาบาล และนำข้อมูลผลลัพธ์ที่ได้จากการทำ Sentiment Analysis มาจัดหมวดวิเคราะห์คำที่ปรากฏในประโยคความคิดเห็นเพื่อให้จัดอยู่ใน Concern Dimension ต่างๆ ให้มีความสอดคล้องกับระบบการจัดการจริงของโรงพยาบาล เพื่อแสดงผลรายงานผ่าน dashboard ต่อผู้บริหารโรงพยาบาล เพื่อให้ทราบถึงปัญหาได้อย่างรวดเร็วและช่วยให้ผู้บริหารตอบสนองผู้ใช้บริการได้อย่างทันถ่วงที

1.2 วัตถุประสงค์ของงานวิจัย

1.2.1 ศึกษาการทำ Sentiment Analysis ด้วยการใช้แบบจำลองการทำนาย (Machine Learning) เพื่อให้ได้ผลลัพธ์ในการทำนายความรู้สึกจากความคิดเห็นของผู้มีประสบการณ์การใช้บริการโรงพยาบาล (Customer Review) ที่มีความถูกต้องและแม่นยำที่สุด

1.2.2 ศึกษาเพื่อหาความรู้เชิงลึกจากการวิเคราะห์ Sentiment และการวิเคราะห์ Concern Dimension เพื่อสร้างการแสดงผลของรายงานในรูปแบบ History Report เพื่อให้ทราบถึงมุมมองในเชิงลึกและปัญหาของผู้ใช้บริการที่มีต่อระบบการจัดการของโรงพยาบาลได้อย่างรวดเร็ว

1.3 ขอบเขตของการวิจัย

1.3.1 งานวิจัยนี้ใช้ข้อมูล External Data จากการเก็บรวบรวมข้อมูลความคิดเห็นของผู้มีประสบการณ์ใช้บริการโรงพยาบาลในประเทศไทย จากเว็บไซต์ www.honestdocs.co.th โดยข้อมูลที่นำมาใช้ในการวิจัย ได้แก่ ชื่อโรงพยาบาล ความคิดเห็นของผู้มีประสบการณ์การใช้บริการโรงพยาบาล คะแนนความพึงพอใจ และวันที่แสดงความคิดเห็น

1.3.2 การสร้างแบบจำลองการเรียนรู้ของเครื่องด้วยเทคนิคอัลกอริทึมที่สนใจ ได้แก่ Multinomial Naïve Bayes, Random Forest และ LSTM เพื่อทำนายประเภทความรู้สึกของข้อความความคิดเห็น

1.3.3 การแปลงข้อความความคิดเห็นภาษาไทยเป็นภาษาอังกฤษก่อน เพื่อนำข้อมูลเข้าสู่กระบวนการเตรียมข้อมูล ด้วยเทคนิคการประมวลผลภาษาธรรมชาติภาษาอังกฤษ ดำเนินการเนื่องด้วยการการใช้การประมวลผลภาษาไทยยังมีให้ใช้งานได้ไม่มากและส่วนใหญ่จะเป็นการใช้ภายในองค์กร

1.4 ประโยชน์ที่คาดว่าจะได้รับจากงานวิจัย

1.4.1 เข้าใจการหลักการทำ Sentiment Analysis ผ่านการเรียนรู้ของเครื่องในการทำนายความรู้สึกของข้อความการแสดงความคิดเห็น

1.4.2 สามารถนำไปประยุกต์ใช้ในการสร้างระบบ Support system เพื่อใช้ดูรายงานผลการแสดงวิเคราะห์ข้อมูลประสิทธิภาพการให้บริการผ่านแดชบอร์ด เพื่อนำข้อมูลเหล่านี้รายงานผลต่อผู้บริหารโรงพยาบาลและผู้กำหนดนโยบาย เพื่อให้ถึงทราบปัญหาในด้านการจัดการคุณภาพเกิดขึ้น และนำสู่การปรับปรุงแก้ไขปัญหาที่เกิดขึ้นได้อย่างรวดเร็วเมื่อมีความคิดเห็นเชิงลบเกี่ยวกับโรงพยาบาลเกิดขึ้น

1.4.3 สามารถนำไปประยุกต์การสร้างระบบการรวิวโรงพยาบาลเพื่อให้ผู้ใช้บริการสามารถเข้ามาแสดงความคิดเห็น ตรวจสอบและคัดเลือกโรงพยาบาลที่มีจุดแข็งและจุดอ่อนได้ตรงตามความต้องการ

บทที่ 2

เอกสารทฤษฎีและงานวิจัยที่เกี่ยวข้อง

2.1 งานวิจัยที่เกี่ยวข้อง

การวิเคราะห์ความรู้สึกจากข้อความการแสดงความคิดเห็นจากประสบการณ์ของผู้ใช้บริการ มีการใช้อย่างกว้างขวางกับเสียงของผู้ใช้บริการ(Voice of customer) เช่น บทวิจารณ์บนโซเชียลมีเดีย และการตอบแบบสำรวจในรูปแบบกระดาษหรือบนสื่อออนไลน์ และได้มีการนำไปใช้ในธุรกิจที่หลากหลายตั้งแต่การทำตลาด และการบริการลูกค้าไปจนถึงทางการแพทย์ ซึ่งในงานวิจัยนี้มุ่งเน้นการนำเสนอการวิเคราะห์ข้อมูลความคิดเห็นของผู้ใช้บริการโรงพยาบาล เพื่อกำหนดทัศนคติของผู้วิจารณ์เกี่ยวกับประสบการณ์การใช้บริการตั้งแต่อดีตจนถึงปัจจุบันที่มีผลต่อการให้บริการรวมถึงคุณภาพของผลิตภัณฑ์ในธุรกิจต่างๆ เพื่อประเมินทัศนคติของบุคคลผ่านข้อความความคิดเห็น โดยการสร้างแบบจำลองการเรียนรู้ของเครื่องเพื่อใช้ในการทำนายประเภทความรู้สึกของข้อความความคิดเห็น ด้วยอัลกอริทึม Multinomial Naive Bayes, Random Forest และ LSTM เพื่อเปรียบเทียบประสิทธิภาพของแบบจำลองการเรียนรู้ด้วยอัลกอริทึมทั้งสามแบบ และนำไปสู่การวิเคราะห์เพื่อค้นหาข้อมูลเชิงลึก

งานวิจัยทางด้านการวิเคราะห์ความรู้สึกได้ถูกใช้มานานและส่วนใหญ่ใช้ในงานการวิเคราะห์ข้อมูลรีวิวในธุรกิจสำหรับการรีวิวผลิตภัณฑ์ ซึ่งมีจำนวนมากในหลากหลายบริบท และในปัจจุบันการแสดงความคิดเห็นและการสื่อสารเพื่อแสดงอารมณ์และความรู้สึกมีมากมายในโซเชียลมีเดียจากแพลตฟอร์มต่างๆ ข้อมูลเหล่านี้จะถูกถ่ายทอดออกมาเป็นข้อมูลเชิงลึกเกี่ยวกับอารมณ์และความรู้สึกของบุคคลที่มีผลทางธุรกิจต่างๆ ในก่อนหน้านี้มีงานวิจัยหลายงานได้แก่(Sajeevan & L. K, 2019)การวิเคราะห์บทความวิจารณ์ภาพยนตร์โดยใช้เทคนิค การเรียนรู้เชิงลึก LSTM และ CNN ร่วมกันเพื่อจัดประเภทบทวิจารณ์ภาพยนตร์ของ IMDB (Sewalk, Tuli, Hswen, Brownstein, & Hawkins, 2018)การวิเคราะห์ข้อมูลของทวิตเตอร์เพื่อทำความเข้าใจบทสนทนาบนเว็บที่เกี่ยวข้องกับประสบการณ์การของผู้ป่วยในพื้นที่ และเวลาที่ต่างกันของผู้ใช้บริการในด้านการรับการรักษาในสหรัฐอเมริกา โดยพัฒนาโมเดล Support Vector Machine ซึ่งเป็นโมเดลที่ให้ประสิทธิภาพในการแยกประเภทที่ดีที่สุด (Jianqiang, Xiaolin, & Xuejun, 2018)การวิเคราะห์ความรู้สึกบนข้อมูลในทวิตเตอร์โดยใช้ Deep learning Neural Network ด้วยอัลกอริทึม Convolutional เพื่อจำแนกประเภทความรู้สึกของทวิต (Hassan et al., 2020)การวิเคราะห์ความรู้สึกเกี่ยวกับบทความที่ถูกพูดถึงในทวิตเตอร์เพื่อวัดผลกระทบของบทความวิจัยผ่านการแสดงความรู้สึกในทวิตเตอร์ โดยสรุปได้ว่าการเพิ่มความคิดเห็นใหม่ๆ ใส่ไว้ในพจนานุกรมเกี่ยวกับ

โดเมนทางวิทยาศาสตร์ในแบบจำลองการถดถอยจะสามารถปรับปรุงค่า R-Square ได้ดียิ่งขึ้น (Haruechaiyasak, Kongthon, Palingoon, & Trakultaweekoon, 2013) การวิเคราะห์ความรู้สึกบนโซเชียลโดยใช้ S-Sense เฟรมเวิร์กสำหรับโดยวิเคราะห์ 2 โมดูลหลัก คือการกำหนดระดับความตั้งใจและระดับความรู้สึกให้เหมาะสมกับข้อความ โดย S-Sense ด้วยการใชแบบจำลองการจำแนกประเภท เพื่อนำไปพัฒนาใช้กับแอปพลิเคชันจำนวนมากรวมถึงการติดตามข้อมูลในโซเชียลมีเดียเพื่อการปรับปรุงแผนการตลาดและการคิดแคมเปญ (Wicaksono & Mariyah, 2019) การวิเคราะห์ความรู้สึกผู้บริโภคในรูปแบบการแสดงความคิดเห็นบนโซเชียลมีเดียจากเพจเฟซบุ๊ก เพื่อประเมินและเปรียบเทียบการวิเคราะห์ค่าความเชื่อมั่นเพื่อทดสอบและเปรียบเทียบวิธีการใช้พจนานุกรมและการเรียนรู้ของเครื่องกับการวิเคราะห์ความรู้สึก (Shah, Isah, & Zulkernine, 2018) การค้นหาความเคลื่อนไหวของราคาหุ้นที่มีความสัมพันธ์กับความคิดเห็น (ความรู้สึก) ของสาธารณชนที่มีต่อบริษัทนั้น อัลกอริทึมที่ใช้พิจารณาความรู้สึกความคิดเห็น ข่าว และราคาหุ้นในอดีตเพื่อคาดการณ์ราคาหุ้นในอนาคต ดำเนินการโดยใช้การเรียนรู้ของเครื่องด้วยโมเดล Deep-learning , Support Vector Machine, MNB, linear regression, Naïve Bayes and LSTM (Long Short-Term Memory) ซึ่งให้ผลลัพธ์ที่สามารถยืนยันความสำเร็จได้

เมื่อหากเราพิจารณาโรงพยาบาล ถือว่าเป็นธุรกิจอย่างหนึ่ง ซึ่งในงานนี้ได้นำข้อมูลการรีวิวประสบการณ์การใช้บริการของผู้ป่วย ซึ่งข้อมูลเหล่านี้มีผลต่อการตัดสินใจของผู้มาใช้บริการโรงพยาบาล ซึ่งงานวิจัยในการวิเคราะห์ความรู้สึกของผู้ใช้บริการโรงพยาบาลมีน้อยมากในประเทศไทย จากการสืบค้นจากฐานข้อมูล Scholar ที่ได้ทำการสืบค้นในวันที่ 14/6/2020 ด้วยคำสำคัญที่ค้นหา คือ Thai, hospital, reviews, sentiment analysis พบว่างานวิจัยในฐานข้อมูลดังกล่าว และตรวจสอบข้อมูลงานวิจัยจำนวน งานวิจัย 100 อันดับแรกที่ค้นหา ไม่พบงานวิจัยที่ใช้คำสำคัญในการสืบ ที่เกี่ยวกับการวิเคราะห์ความรู้สึกการรีวิวที่เกี่ยวข้องกับด้านการบริการของโรงพยาบาล ด้านสุขภาพ จึงสามารถสรุปได้ว่ายังไม่มี การนำเทคนิคการเรียนรู้ของเครื่องมาใช้ในการวิเคราะห์ข้อมูลการรีวิวประสบการณ์การใช้บริการของผู้ป่วยในประเทศไทยจึงเป็นเหตุผลที่ทำงานวิจัยนี้ขึ้นมาเพื่อตอบสนองความต้องการในการปรับปรุงธุรกิจให้ดีขึ้น ในงานวิเคราะห์ความรู้สึกข้อมูลประเภทข้อความ วิธีการส่วนใหญ่ใช้แบบจำลองการเรียนรู้ด้วยอัลกอริทึม multinomial Naive Bayes ซึ่งเป็นอัลกอริทึมพื้นฐานในการแยกประเภทความรู้สึกของข้อความ เนื่องจากสามารถใช้งานง่าย ประมวลผลเร็ว แต่ก็ยังมีข้อจำกัดในที่เป็นลักษณะการสร้างคลังคำศัพท์ โดยอาจจะยังทำงานได้ไม่ดีกับข้อมูลที่มีความสัมพันธ์กันเชิงลำดับ ซึ่งในขณะเดียวกัน อัลกอริทึมที่มีประสิทธิภาพในการเรียนรู้ที่ซับซ้อนอย่างการเรียนรู้เชิงลึกในรูปแบบ LSTM การ

ทำงานของอัลกอริทึมนี้จะสามารถทำงานได้ดีกับข้อมูลที่มีความสัมพันธ์กันเชิงลำดับ แต่ยังมีข้อจำกัดคือ ต้องการข้อมูลจำนวนมากในการสร้างแบบจำลองและใช้เวลาค่อนข้างนาน นั่นคือ LSTM อาจจะไม่สามารถทำงานได้ดีเสมอไปกับทุกชุดข้อมูล ในงานวิจัยนี้เราจึงสนใจโมเดลเหล่านี้ว่า ถ้านำมาทดสอบประสิทธิภาพบนชุดข้อมูลการรวิวโรงพยาบาลว่าจะมีพฤติกรรมอย่างไร ในการทดลองได้ใช้โมเดล Naïve Bayes, LTSM และ Random Forest มาประเมินและเปรียบเทียบประสิทธิภาพการในแยกประเภทความรู้สึกรู้สึกของข้อมูลชุดนี้

2.2 เอกสารทฤษฎีที่เกี่ยวข้อง

2.2.1 การวิเคราะห์ความรู้สึก

การวิเคราะห์ข้อมูลในปัจจุบันไม่ได้มีการมุ่งเน้นไปที่ตัวเลขในการประเมินเพียงอย่างเดียวอีกต่อไป เมื่อมีข้อมูลที่หลากหลายมากขึ้น เช่น ข้อมูลที่เป็นข้อความ รูปภาพ เสียง ก็เป็นข้อมูลที่สำคัญที่ให้ข้อมูลเชิงลึกที่มีค่าด้วยเช่นกัน โดยในงานวิจัยนี้จะมุ่งเน้นไปที่การนำข้อมูลข้อความที่สามารถรวบรวมได้จากการแสดงความคิดเห็นของผู้มีประสบการณ์ใช้บริการโรงพยาบาลมาใช้ในการวิเคราะห์ ซึ่งการดึงและรวบรวมข้อมูลดังกล่าวจะพูดถึงในหัวข้อถัดไป การวิเคราะห์ความรู้สึกของประชากรถือเป็นแอปพลิเคชันหนึ่งในการวิเคราะห์ข้อความสำคัญในภาษามนุษยชาติที่สามารถนำไปใช้ประโยชน์ในการพัฒนาธุรกิจ ในการวิเคราะห์นี้เป็นการจำแนกประเภทแบบไบนารี ซึ่งผลลัพธ์จะออกมาในสองรูปแบบคือ Positive และ Negative ด้วยการใช้แบบจำลองการเรียนรู้ของเครื่อง(Machine Learning) ด้วยการเรียนรู้ข้อมูลจากข้อมูลชุดฝึกฝนที่ถูกต้องจากนั้นข้อมูลที่ถูกเรียนรู้จะถูกนำไปใช้ในการทำนายกับชุดข้อมูลทดสอบ ตัวอย่างการวิเคราะห์ความรู้สึกในโลกปัจจุบัน เช่น การตรวจสอบความรู้สึกของผู้บริโภคกับแบรนด์และผลิตภัณฑ์ของธุรกิจ โดยการมีนักวิจัยรวบรวมผลการวิเคราะห์เหล่านั้นเพื่อนำมาทำความเข้าใจมุมมองของประชากรที่มีต่อสิ่งใดสิ่งหนึ่ง เพื่อใช้ในการปรับปรุงผลิตภัณฑ์จากการรวบรวมข้อมูลเพื่อให้สามารถระบุได้ว่าสิ่งที่ผู้บริโภคไม่พอใจหรือพึงพอใจอะไร

2.2.2 การเก็บรวบรวมข้อมูลจากเว็บไซต์

Web scraping คือการเก็บรวบรวมข้อมูลหรือการสกัดข้อมูลเพื่อใช้สำหรับการดึงข้อมูลจากเว็บไซต์ โดยใช้ไลบรารี selenium ในการเข้าถึงเว็บไซต์ผ่าน web driver ในระหว่างการสกัดข้อมูลจะสามารถสร้างระบบการทำงานแบบอัตโนมัติ จากการใช้ selenium เพื่อเปิดเว็บไซต์และเก็บรวบรวมข้อมูลได้ ซึ่งเป็นหนึ่งในวิธีการคัดลอกข้อมูลจากเว็บไซต์โดยข้อมูลจะถูกเก็บในรูปแบบไฟล์ที่เป็นตาราง เช่น excel, csv, sql, text เพื่อใช้ในการค้นหาในภายหลัง เนื่องจากข้อมูลหรือเนื้อหาต่างๆส่วนใหญ่บนเว็บไซต์ เข้าถึงได้ด้วยการใช้เว็บเบราว์เซอร์ ซึ่งเว็บเบราว์เซอร์ไม่มีฟังก์ชันที่ใช้

สำหรับคัดลอกหรือดาวน์โหลดข้อมูลที่จะมาใช้ได้โดยตรง วิธีการที่จะได้ข้อมูลมาได้คือการใช้ระบบ manual เมื่อข้อมูลมีมากขึ้นจะใช้เวลานาน ดังนั้นการทำ Web scraping จะมีประโยชน์ในการทำงานด้วยระบบอัตโนมัติเมื่อเราต้องการดึงข้อมูลออกมาจากเว็บไซต์ที่เราต้องการ

2.2.3 การสกัดข้อมูลจากเว็บไซต์

1. การดึงและสกัดข้อมูลเพื่อเก็บรวบรวมข้อมูลในเว็บไซต์ สิ่งแรกในการทำคือ การเปิดหน้าเว็บเบราว์เซอร์ผ่านไลบรารี selenium และส่งคำขอเว็บเพจที่ต้องการเปิดไปยังเว็บเซิร์ฟเวอร์ด้วยการใช้ไลบรารี Python request ซึ่งจะสามารถเก็บเนื้อหา HTML code ของหน้าที่ส่งคำขอไปได้

2. การสกัดเนื้อหาในส่วนที่เราต้องการโดยใช้ไลบรารี BeautifulSoup เพื่อใช้ในการระบุ element ต่างๆที่ต้องการเก็บรวบรวมจากเนื้อหาของ HTML นั้น

เนื่องจากในกระบวนการทำงานในงานวิจัยนี้ต้องการเก็บรวบรวมข้อมูลความคิดเห็นของทุกโรงพยาบาลในเว็บไซต์ www.honestdoc.co.th จึงใช้กระบวนการทำงานแบบซ้ำ เพื่อเปิดเว็บเบราว์เซอร์สำหรับทุกโรงพยาบาลและทำการสกัดข้อมูลเพื่อดึงข้อมูลมาเก็บไว้ในรูปแบบ Data frame และนำข้อมูลเข้าไปเก็บไว้ในฐานข้อมูลในรูปแบบไฟล์ .csv เพื่อนำข้อมูลไปใช้ในขั้นตอนการเตรียมชุดข้อมูลเพื่อนำใส่เข้าโมเดลต่อไปซึ่งจะอธิบายไว้ในหัวข้อถัดไป

2.2.4 การประมวลผลภาษาธรรมชาติ

การประมวลผลภาษาทางธรรมชาติ (Natural language processing) คือโมดูล NLTK (Natural Language Toolkit) ซึ่งเป็นชุดโปรแกรมสำหรับการประมวลผลภาษาธรรมชาติ(NLP) สำหรับภาษาอังกฤษ เป็นโมดูลในภาษาไพทอนและเป็นที่ยอมรับกันทั่วโลกซึ่งจะใช้เทคนิคนี้ ในการเตรียมข้อมูลสำหรับนำเข้าโมเดลเพื่อให้เครื่องฝึกฝนและทดสอบกับชุดข้อมูลเพื่อให้โมเดลเข้าใจภาษาของข้อความในความคิดเห็นของผู้ใช้บริการมากที่สุด การประมวลผลภาษาธรรมชาติเป็นระบบที่ทำให้คอมพิวเตอร์เข้าใจภาษามนุษย์เพื่อให้คอมพิวเตอร์ประมวลเป็นคำสั่งที่เป็นภาษาในชีวิตประจำวันที่คอมพิวเตอร์นำไปใช้งานได้ เช่น นาฬิกาปลุกพูดได้ เครื่องคิดเลขพูดได้ ระบบตรวจจับคำผิด เมื่อมีผู้เขียนบทวิจารณ์หรือการเขียนแสดงความคิดเห็น ซึ่งสามารถเป็นภาษาที่ไม่เป็นทางการ ดังนั้นจะทำให้เครื่องเข้าใจคำในแบบที่แตกต่างกันออกไปถ้าไม่มีการสอนซึ่งถือว่าเป็นปัญหาในการวิเคราะห์ข้อความที่เกิดขึ้นรวมถึงปัญหาของ การสะกดคำผิด ตัวอักษรใหญ่ จึงต้องจัดการปัญหาเหล่านี้ด้วยการประมวลผลข้อความและเตรียมข้อมูลข้อความไว้สำหรับการนำข้อมูลเพื่อสร้างแบบจำลองการเรียนรู้ของเครื่อง

2.2.5 การสร้างคุณลักษณะของข้อมูล

การใช้เทคนิคการเรียนรู้ของเครื่อง(Machine Learning) เพื่อใช้วิเคราะห์ความรู้สึก จำเป็นต้องเปลี่ยนข้อความเหล่านี้ให้เป็นค่าตัวเลขซึ่งถูกเรียกว่า ตัวแทนเวกเตอร์ สำหรับคำ ซึ่งสามารถทำได้หลายเทคนิคในงานวิจัยนี้

ในขั้นตอนการคำนวณความถี่ จะนำคำไม่ซ้ำ(unique word) มาคำนวณหาความถี่ที่ปรากฏ(TF) ในทุกประโยคดังตารางที่ 1 ซึ่งค่าที่ได้จาก TF จะนำไปคำนวณค่าความผกผันของจำนวนคำที่ปรากฏ(IDF) ดังตารางที่ 2 จากนั้นจะเป็นการคำนวณค่า TF-IDF โดยการคำนวณจากตารางที่ 3

ตาราง 1 ตารางตัวอย่างการคำนวณความถี่ของคำที่ปรากฏ(TF Calculation)

features	term	TF(term, Doc1)	TF(term, Doc2)	...
1	doctor	1	0	
2	treat	0	1	
3	time	1	1	
...	...	..	..	
8300	bed	0	1	

ตาราง 2 ตารางตัวอย่างการคำนวณค่าความถี่ผกผัน(IDF Calculation)

features	term	TF(term, Doc1)	TF(term, Doc2)	IDF(term, doc1)	IDF(term, doc2)
1	doctor	1	0	1.41	2.10
2	treat	0	1	2.10	1.41
3	time	1	0	1.41	2.10
...	...	..	..	...	...
8300	bed	0	1	2.10	1.41

ตาราง 3 ตารางตัวอย่างการคำนวณTFIDF

features	term	TF-IDF(term,Doc1)	TF-IDF(term,Doc1)
1	doctor	1.41	0
2	treat	0	1.41
3	time	1.41	0
...	...	..	..
8300	bed	0	1.41

2.2.5.1 การสร้างคุณลักษณะด้วยเทคนิค TFIDF

เทคนิคการคำนวณ ความสำคัญของคำในเอกสาร(TF-IDF) TF คือ Term Frequency การนับความถี่ของคำ (t) ในเอกสาร (d)

$$T_{dt}, d = \text{count}(t, d) \quad (1)$$

หรืออีกวิธีหนึ่งเพื่อปรับสเกลไม่ให้มีขนาดที่ใหญ่เกินไป (+1 ใส่เพื่อป้องกันไม่ให้เกิดเคส log 0)

$$T_{ft}, d = \log_{10}(\text{count}(t, d) + 1) \quad (2)$$

IDF คือ Inverse Document Frequency ในการทำคำนวณ IDF จำเป็นต้องรู้จัก DF คือ Document Frequency เป้าหมายของการทำ DF คือ จะให้ weight เยอะกับคำที่พบบ่อยในเอกสารซึ่งจะถือว่าเป็นคำที่สำคัญ แต่ถ้าคำนั้นพบบ่อยในเอกสารแสดงว่าคำนั้นไม่ค่อยมีความสำคัญ ซึ่งเราจะใช้ Inverse Document Frequency แทนการตีความหมายได้ง่ายกว่า DF ดังสูตรคือ

$$\frac{N}{d_{ft}} \quad (3)$$

N คือจำนวน document ใน collection ส่วน df คือจำนวน document ที่มีคำนั้นๆ และมีการ take log ฐาน 10 เข้าไปเพื่อการปรับสเกล.

$$idf = \log_{10}\left(\frac{N}{d_{ft}}\right) \quad (4)$$

ดังนั้นเมื่อค่า IDF ยังมีค่าสูงมากๆ คำๆนั้นยังมีความสำคัญมาก

TF-IDF คือค่ามาตรฐานสำหรับการพิจารณาค่า weight ของ co-occurrence matrix จะใช้เพื่อหาความเหมือนกันของคำ (word similarity) วิธีการหาค่าของ TFIDF คือการนำค่า TF และ

IDF มาคูณกันเพื่อให้ weight ของ TF และ weight ของ IDF ที่ซึ่งมีค่าตรงข้ามกันนั้นเป็นน้ำหนักที่สามารถแยกคำสำคัญออกมาได้จากสูตรด้านล่าง ซึ่งประโยชน์จากการหาค่า TFIDF คือ การแยกคำที่ไม่ต้องการออกมา โดยทำการฟิลเตอร์ คำที่มีน้ำหนักต่ำออก และยังสามารถนำไปใช้ประโยชน์ในการทำ document similarity โดยการคำนวณ tf-idf ของทุกคำใน document แล้วนำมาหา centroid แล้วนำมาเปรียบเทียบกับ $\cos(d1, d2)$ ในงานวิจัยนี้ได้ใช้เทคนิค TF-IDF จากการใช้ไลบรารี Tfidf Vectorizer ของ SKlearn ดังสมการคำนวณด้านล่าง

$$Tfidf(i,j) = tf_{i,j} \times \log(N/df_i) \quad (5)$$

$Tf_{i,j}$ = จำนวนครั้งทั้งหมดของ i ใน j

df_i = จำนวนเอกสารทั้งหมด (speeches) ที่มี i

N = จำนวนเอกสารทั้งหมด (speeches)

ที่มา (Li, 2020)

2.2.5.2 การสร้างคุณลักษณะด้วยเทคนิค Word Embedding.

Word Embedding คือการคำนวณค่าความคล้ายคลึงของคำอื่นๆในบริบทของคำที่แตกต่างกัน ด้วยการสร้าง เวกเตอร์คุณลักษณะ (Feature Vector) จาก token ของคำทุกคำที่อยู่ในชุดข้อมูลโดยการแปลงค่าเป็นตัวเลขที่สามารถนำไปใช้คำนวณต่อได้ในอัลกอริทึมการคำนวณค่าความคล้ายคลึงกันของ Word Embedding วิธีการสร้างเวกเตอร์คุณลักษณะจะเริ่มด้วยการเข้ารหัส One-Hot Encoding ของคำแต่ละคำ เพื่อให้ข้อมูลอยู่ในรูปแบบที่คอมพิวเตอร์สามารถทำความเข้าใจได้ ซึ่งการทำงานของ One-Hot Encoding จะทำการนำจำนวนคำที่ปรากฏในชุดข้อมูล และประโยค ทำเป็นเวกเตอร์ เท่ากับจำนวนคำที่ปรากฏบนชุดข้อมูลจากนั้นทำการรวมเวกเตอร์ที่ได้จาก One Hot Encoding ซึ่งสามารถกำหนดขนาดของ Dimension ได้

cat =>	1.2	-0.1	4.3	3.2
mat =>	0.4	2.5	-0.9	0.5
on =>	2.1	0.3	0.1	0.4

ภาพประกอบ 1 แนวคิดการสร้าง One Hot Encoding

ที่มา https://www.tensorflow.org/text/guide/word_embeddings

การพัฒนาการเรียนรู้ค่าความสอดคล้องของ Word Embedding ในปัจจุบันมีการพัฒนาอัลกอริทึมต่างๆขึ้นมา แต่ที่เป็นที่นิยมในการใช้งาน คือ Word2Vec

Word2Vec คืออัลกอริทึมที่ถูกพัฒนาโดย Google ที่จะมีการเรียนรู้ความสัมพันธ์ของคำด้วย โครงสร้างให้มีโครงข่ายประสาทเทียม 2 ชั้น มาใช้ในการฝึกสอน โดยอาศัยสมมติฐานการกระจายจากสมมติฐานการกระจายระบุว่าคำที่มีค่าใกล้เคียงเหมือนกันมักจะมีความหมายคล้ายกัน วิธีนี้ช่วยจับคู่คำที่มีความหมายคล้ายกันกับ embedding vectors

2.2.6 เทคนิคการเรียนรู้ของเครื่อง

การศึกษาอัลกอริทึมของคอมพิวเตอร์ซึ่งถือเป็นการเรียนรู้ของเครื่องซึ่งเป็นส่วนหนึ่งของปัญญาประดิษฐ์ ในการทำงานของอัลกอริทึมจะสร้างแบบจำลองทางคณิตศาสตร์ด้วยข้อมูลตัวอย่างที่ใช้ในการสอน เพื่อใช้ช่วยในการตัดสินใจ และในการสร้างอัลกอริทึมจะสามารถเรียนรู้ข้อมูลเพื่อทำนายข้อมูลได้ แทนการทำงานตามลำดับคำสั่งของโปรแกรมคอมพิวเตอร์ในการเรียนรู้ของเครื่องมีความเกี่ยวข้องอย่างมากในสาขาสถิติ ประเภทหลักๆของ machine learning มีดังนี้ Supervised Learning เป็นการเรียนรู้ของเครื่องโดยมี data มาให้เครื่องเรียนรู้ Unsupervised Learning เป็นการเรียนรู้ของเครื่องโดยไม่มี data มาให้เครื่องเรียนรู้ และ Reinforcement Learning คือ การเรียนรู้ตามสภาพแวดล้อมผ่านการให้รางวัล

2.2.7 อัลกอริทึม Multinomial Naïve Bayes

Naïve Bayes เป็นเทคนิคในการสร้างตัวแยกประเภทโดยใช้ความน่าจะเป็นเข้ามาคำนวณ แบบจำลองในการกำหนดป้ายกำกับ ซึ่งเป็นอัลกอริทึมที่ไม่ซับซ้อน การเรียนรู้ของเครื่องจักรด้วยนาอิวเบย์สามารถแบ่งออกได้เป็น 2 แบบจำลอง คือ Multivariate Bernoulli และ Multinomial

$$P(c|x) = \frac{P(x|c)P(c)}{P(x)} \quad (6)$$

$$P(c|x) = P(x_1|c) \times P(x_2|c) \times \dots \times P(x_n|c) \times P(c) \quad (7)$$

อธิบายสมการ แทนตัวแปร 3 ตัว คือ

C คือ คลาส (Class) x คือ แอตริบิวต์ (Attribute)

P คือ Probability (ความน่าจะเป็น)

P(c|x) Posterior probability คือ ความน่าจะเป็นที่ข้อมูลที่มีแอตทริบิวต์ x จะมีคลาส C

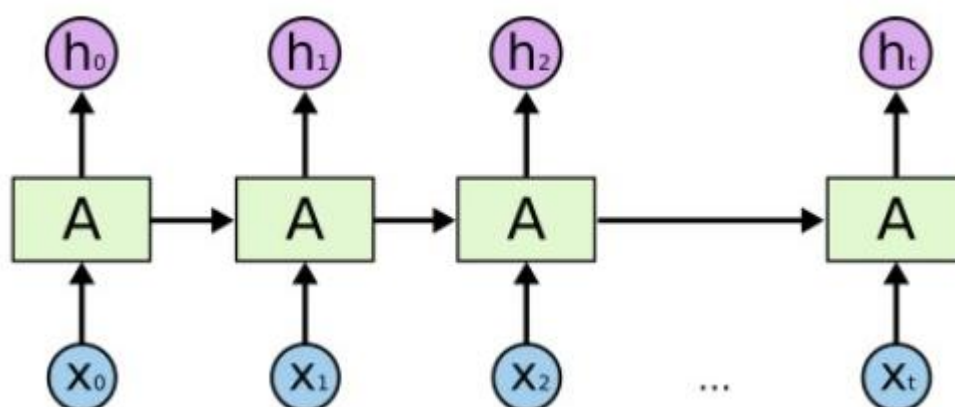
P(x|c) Likelihood คือ ความน่าจะเป็นที่ข้อมูลที่มีคลาส C และมีแอตทริบิวต์ x

$P(c)$ Prior probability คือ จำนวน Class ที่อาจจะเกิดขึ้นจำนวน Class ทั้งหมด หรือความน่าจะเป็นของ Class C

$P(x)$ Predictor Prior probability คือ จำนวน Attribute ทั้งหมด

2.2.8 อัลกอริทึม Long-short-term memory

Long Short-Term Memory เป็นโครงข่ายประสาทเทียมแบบหนึ่งที่ถูกออกแบบมา สำหรับการประมวลผลลำดับ (sequence) LSTM จัดเป็นโครงข่ายประเภท Recurrent Neural Network (RNN) คือ Neural network ที่มีการนำเอา output ของการคำนวณก่อนหน้านี้กลับมาใช้ ใหม่ เช่นดังรูปด้านล่าง โดยข่ายนั้นคือตัวอย่าง RNN ที่มี 1 layer และ 1 node ใน layer นั้น เมื่อนำ RNN ไปใช้ประมวลผลลำดับ $(x(1), y(1)), \dots, (x(10), y(10))$ $x(t)$ คือข้อมูลนำเข้า ณ เวลา t และ $y(t)$ คือค่าส่งออกที่ต้องการ เช่น $x(1), \dots, x(10)$ ซึ่งคือลำดับคำในภาษาไทย ส่วน $y(1), \dots, y(10)$ เป็น part of speech ของคำเหล่านี้ซึ่งสามารถนำไปใช้งาน sequence-to-sequence อื่นได้

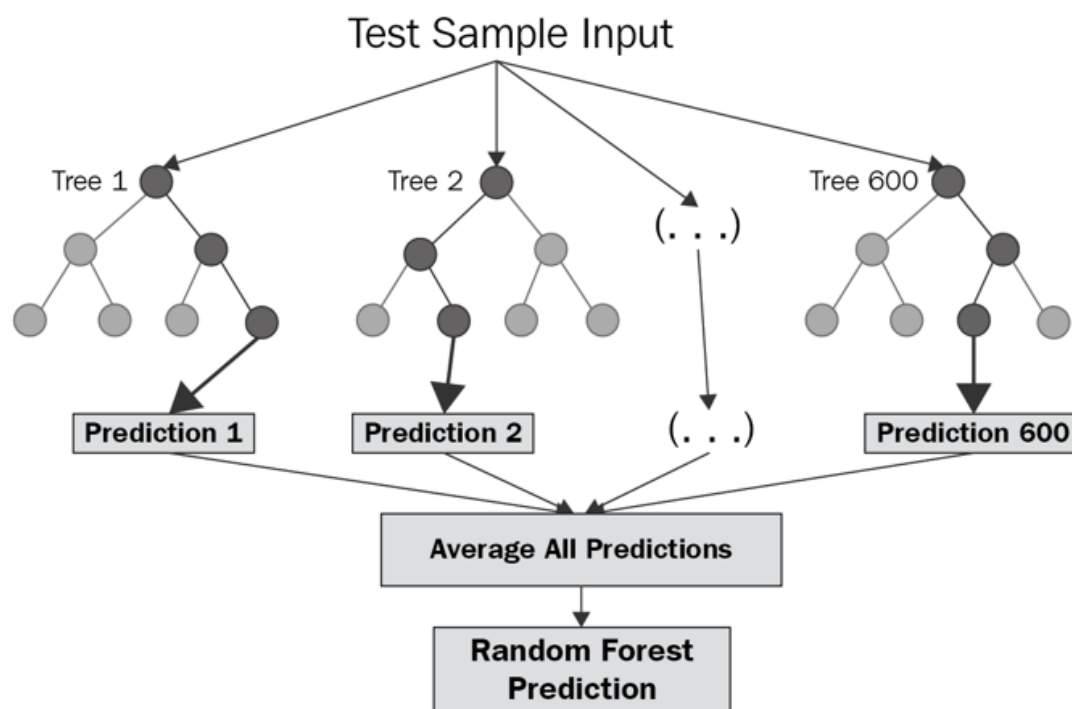


ภาพประกอบ 2 หลักการสร้างแบบจำลองการเรียนรู้ด้วยอัลกอริทึม LSTM

2.2.9 อัลกอริทึม Random Forest

Random forest เป็นหนึ่งในอัลกอริทึมการเรียนรู้ของเครื่องที่สามารถนำมาใช้งานใน ปัญหา regression และ classification และเป็นหนึ่งในกลุ่มโมเดลที่เรียกว่า Ensemble learning โดย หลักการของอัลกอริทึม Random Forest คือ สร้าง model ด้วยการรวมการทำงาน ของ อัลกอริทึม Decision Tree หลายๆ model โดยแต่ละ model จะได้รับ ข้อมูลที่ไม่เหมือนกันในการ ฝึกฝน ซึ่งข้อมูลที่ถูกแบ่งไปทำงานในแต่ละโมเดลจะเป็น subset ของชุดข้อมูลทั้งหมด ซึ่งใน ขั้นตอนการทำนาย ก็ให้แต่ละ โมเดล Decision Tree ทำการทำนาย ของใครของมัน และคำนวณ ผลการทำนายสุดท้าย ด้วยการ vote เลือกผลลัพธ์ที่ถูกเลือกโดย Decision Tree ที่มีการทำนาย

เอาพุตที่นั้นมากที่สุด (กรณี classification) หรือ หาค่า mean จาก เอาพุตที่ ของแต่ละโมเดล Decision Tree (กรณี regression) ซึ่ง Decision Tree แต่ละ model ใน Random Forest เป็นการเรียนรู้แบบ weak learner แต่เมื่อนำแต่ละ Decision Tree มาทำนายผลลัพธ์ ร่วมกัน ซึ่งก็จะได้โมเดล Random Forest ที่มีความฉลาด และแม่นยำมากกว่า Decision Tree ที่ทำการทำนายแบบโมเดลเดี่ยว รูป ซึ่งหลักการนี้การเรียนรู้ข้อมูลจะต้องเรียนรู้อย่างเป็นอิสระต่อกันมากที่สุด



ภาพประกอบ 3 หลักการสร้างแบบจำลองการเรียนรู้ด้วยอัลกอริทึม Random Forest

ที่มา <https://corporatefinanceinstitute.com/resources/knowledge/other/random-forest/>

ขั้นตอน Ensemble Method สำหรับการ Voting Bagging หรือ Bootstrap Aggregation เป็นวิธีการรวมกลุ่มที่เรียบง่ายและมีประสิทธิภาพมาก มีขั้นตอนทั่วไปที่สามารถใช้ลดความแปรปรวนของอัลกอริทึมที่มีความแปรปรวนสูงได้ ซึ่งจะถูกใช้ในโมเดล Random forest ที่ประกอบไปด้วย Decision trees มีความอ่อนไหวต่อข้อมูลโดยเฉพาะข้อมูลที่ใช้ในการฝึกฝน หากข้อมูลการฝึกฝนมีการเปลี่ยนแปลง โครงสร้างการตัดสินใจที่ได้จะแตกต่างกันมาก และในทางกลับกัน การทำนายผลลัพธ์ที่ได้ก็อาจจะแตกต่างกันมาก

2.2.10 เทคนิคการแสดงผลของคำ

Word cloud เป็นเทคนิคในการทำ Data visualization ซึ่งคำจากข้อความที่จะแสดงเป็นแผนภูมิแสดงความถี่ ในเทคนิคนี้ คำที่ปรากฏขึ้นบ่อยในข้อความจะแสดงในแบบอักษรใหญ่และโดดเด่นในขณะที่คำที่ปรากฏไม่บ่อยจะแสดงตัวอักษรแบบเล็กลงหรือบาง ซึ่งมีประโยชน์ในการทำ NLP ทำให้เราสามารถวิเคราะห์ข้อความที่ปรากฏได้อย่างรวดเร็ว ซึ่งในงานวิจัยนี้ใช้เทคนิค Word cloud แสดงคำที่ปรากฏในข้อความหลังจากการวิเคราะห์ความรู้สึกในเชิงบวกและเชิงลบของความคิดเห็น เพื่อใช้ระบุค้ายอดนิยมในข้อความ ซึ่งจะช่วยให้เห็นภาพรวมของคำที่ปรากฏขึ้นว่าผู้ที่มีประสบการณ์การใช้บริการโรงพยาบาลพูดถึงอะไรเป็นส่วนใหญ่ในทางบวกและในทางลบ และทำให้เราสามารถนำข้อมูลเหล่านี้ไปใช้ประโยชน์ต่อการด้านการวิเคราะห์แยกหมวดหมู่ของคำ

2.2.11 แนวคิดในการออกแบบแดชบอร์ด

แนวคิดการสร้างแดชบอร์ดแสดงผลรายงานการวิเคราะห์ในการออกแบบแดชบอร์ดสำหรับงานวิจัยนี้ยึดหลักการในการสร้างแดชบอร์ด 4 ข้อดังนี้

1.ผู้ใช้งาน (Users)

การออกแบบ Dashboard สำหรับบุคลากรในโรงพยาบาลที่มีหน้าที่รับผิดชอบในการกำหนดนโยบายหรือควบคุมคุณภาพโรงพยาบาล โดยการออกแบบนี้ Dashboard นี้จำเป็นต้องช่วยตัดสินใจ ติดตามผลการดำเนินงาน ให้เตือนเมื่อมีสิ่งผิดปกติ ในด้านต่างๆของโรงพยาบาล อาทิเช่น การบริการคนไข้ เป็นต้น

2.เนื้อหา (Content)

ข้อมูลที่ถูกมาใช้ในงานวิจัยนี้เป็นข้อมูลรีวิวกของผู้ใช้บริการ ซึ่งสามารถนำมาจัดกลุ่มในด้านต่างๆ ของโรงพยาบาลที่ถูกกล่าวถึงในการรวิวนั้นๆ ในแต่ละช่วงเวลา ว่ามีจำนวนมากน้อยแค่ไหน ที่เป็นเชิงลบหรือเชิงบวก และจำนวนคำเฉพาะคำใดที่ถูกกล่าวถึงบ่อย

3.การนำเสนอข้อมูล (Presentation)

การนำเสนอข้อมูลบนกราฟ เนื่องจากงานวิจัยนี้จะมุ่งเน้นไปทาง History report จึงมีการใช้งานกราฟเส้นที่สื่อถึงข้อมูลที่มีความต่อเนื่องกันแบบรายเดือน โดยจะแสดงจำนวนข้อมูลในเชิงลบเชิงบวกของด้านต่างๆในแต่ละเดือนและนำเสนอจำนวนคำเฉพาะที่ถูกพูดถึงบ่อยครั้งที่สุดในรูปแบบ words cloud และ Pie Graph ที่ใช้แสดงสัดส่วนของด้านต่างๆที่ปรากฏในชุดข้อมูล

4.รูปแบบทิศทางการนำเสนอ (Navigation)

การนำเสนอ Dashboard ภาพรวมของข้อมูลให้ ทุกกราฟที่แสดงมีความเชื่อมโยงกัน โดยเน้นไปที่การแสดงความแตกต่างในแต่ละหมวดหมู่ที่ถูกกล่าวถึง และ Dashboard สำหรับ การนำเสนอคำที่ถูกกล่าวถึงบ่อย ในแต่ละด้านของโรงพยาบาล เนื่องจากแดชบอร์ดมีความสัมพันธ์กับ

ตัวชี้วัดในหัวข้อของปัญหาที่มีความเกี่ยวข้องกับข้อมูลการรื้อรื้อของประสบการณ์ผู้ใช้บริการโรงพยาบาลโดยในงานวิจัยนี้สนใจหัวข้อปัญหา (Concern Dimension) ในด้านระบบการจัดการโรงพยาบาล 5 ด้าน ประกอบด้วย ด้านการบริการ ด้านเจ้าหน้าที่ ด้านโครงสร้างของโรงพยาบาล ด้านการเงิน ด้านกระบวนการดำเนินงาน และด้านอื่นๆ

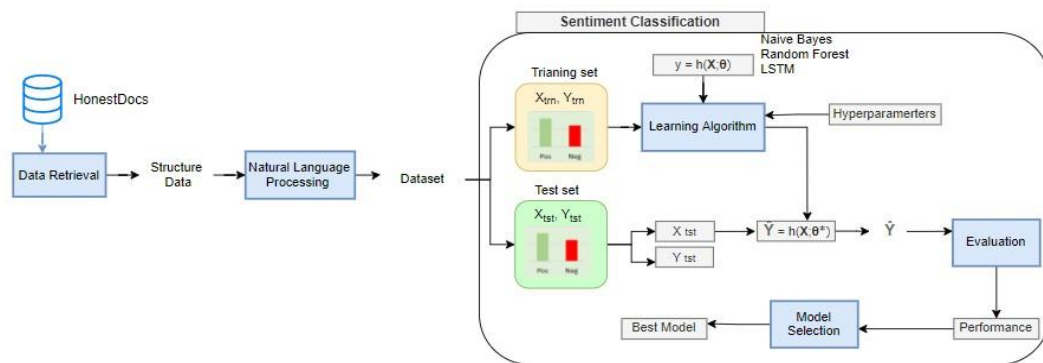


บทที่ 3

วิธีการดำเนินงานวิจัย

สำหรับการดำเนินงานวิจัยนี้ได้นำเสนอการสร้างแบบจำลองการเรียนรู้ของเครื่อง (Machine Learning) เพื่อทำนายความรู้สึก (Sentiment Analysis) จากข้อความแสดงความคิดเห็นจากประสบการณ์ของผู้ใช้บริการโรงพยาบาลและสร้างรายงานการแสดงผลเพื่อวิเคราะห์หาข้อมูลเชิงลึกเพื่อให้ทราบถึงปัญหาที่เกิดขึ้นในระบบการจัดการของโรงพยาบาล ตามหัวข้อตัวชี้วัด (Concern Dimension) ซึ่งมีขั้นตอนและวิธีการดำเนินงานได้ดังนี้

3.1 การวางแผนการดำเนินงานวิจัย



ภาพประกอบ 4 แผนผังการดำเนินงานวิจัย

ในกระบวนการดำเนินงานวิจัย จะเริ่มจากการเก็บรวบรวมข้อมูลรีวิวจากประสบการณ์การใช้บริการโรงพยาบาล จากแหล่งข้อมูลในเว็บไซต์ จากนั้นนำข้อมูลมาผ่านกระบวนการประมวลผลภาษาธรรมชาติเพื่อจัดเตรียมข้อมูล เมื่อข้อมูลอยู่รูปแบบที่พร้อมใช้งานแล้วข้อมูลมาแบ่งเป็นชุดข้อมูลฝึกฝนและชุดข้อมูลทดสอบ เพื่อเตรียมข้อมูลเข้าสู่แบบจำลองการเรียนรู้ของเครื่องเพื่อเรียนรู้และทดสอบประสิทธิภาพของแบบจำลองด้วยอัลกอริทึมต่างๆที่นำมาทดสอบในงานวิจัย จากนั้นนำโมเดลที่มีประสิทธิภาพสูงสุดนำไปประยุกต์ใช้การวิเคราะห์ความรู้สึกในชุดข้อมูลใหม่ต่อไป โดยในการสร้างแบบจำลองการเรียนรู้ของเครื่อง กำหนดให้

$$y = h(X; \theta) \text{ แทนแบบจำลองการเรียนรู้ของเครื่อง}$$

$$\hat{y} = h(X; \theta^*) \text{ แทนแบบจำลองที่ปรับพารามิเตอร์แล้ว}$$

X คือ ชุดข้อมูลฝึกฝนและชุดข้อมูลทดสอบสำหรับแบบจำลองการเรียนรู้

๐ คือ พารามิเตอร์ที่ใช้ในการสร้างแบบจำลอง

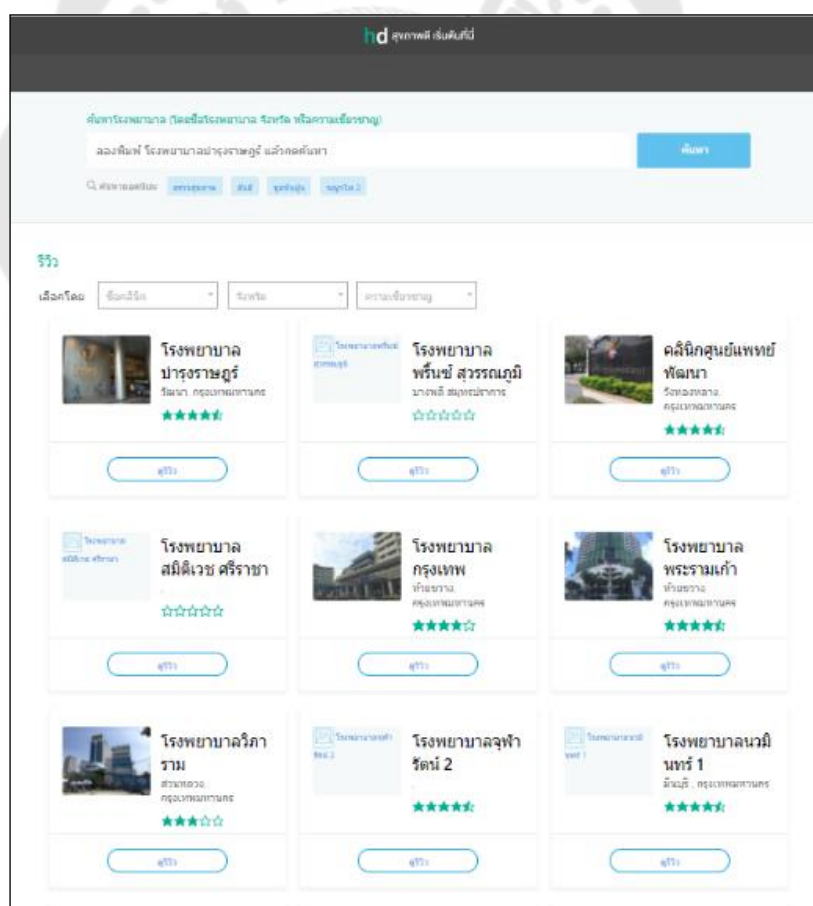
๐* คือ พารามิเตอร์ที่ดีที่สุดที่ได้จากการปรับปรุงประสิทธิภาพแบบจำลอง

h คือ ฟังก์ชันสำหรับการทำงานของแบบจำลอง

y คือ ค่าผลการทำนายที่ต้องการที่ได้จาก

3.2 การเก็บรวบรวมข้อมูล

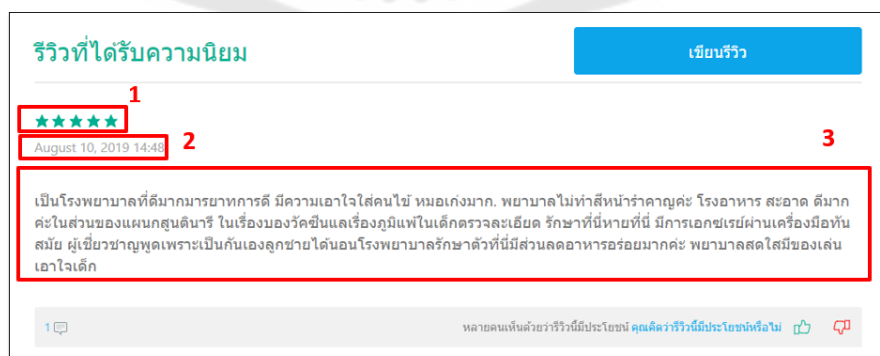
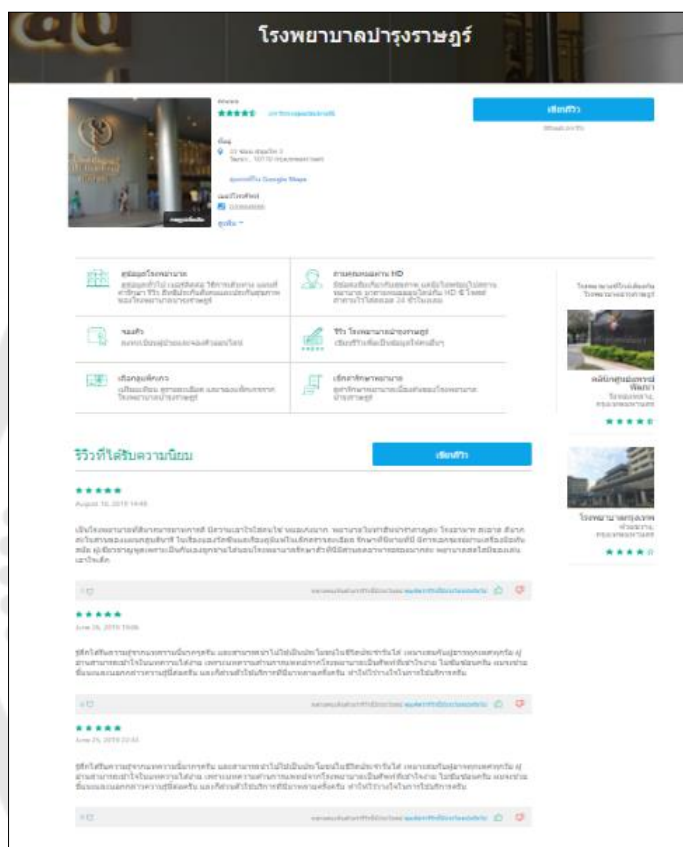
ข้อมูลที่นำมาใช้ในการงานวิจัยนี้คือข้อมูลรีวิวของผู้มีประสบการณ์การใช้บริการโรงพยาบาลทั่วประเทศไทยที่ได้ทำการลงทะเบียนสำหรับเปิดพื้นที่ให้แสดงความคิดเห็นของการเข้าใช้บริการในโรงพยาบาลบน www.honestdocs.com โดยปัจจุบันมีจำนวนโรงพยาบาลที่ได้รับการรีวิวปรากฏในเว็บไซต์จำนวน 369 แห่ง และมีจำนวนรีวิวทั้งหมดรวมทุกโรงพยาบาล 5,485 รีวิว



ภาพประกอบ 5 หน้าเว็บไซต์แสดงรายชื่อโรงพยาบาลที่ได้รับการรีวิว

การเก็บรวบรวมข้อมูลจากเว็บไซต์มีขั้นตอนดังนี้

1. แหล่งข้อมูลที่ต้องการเก็บรวบรวมจะปรากฏอยู่ในหน้าเว็บไซต์การรีวิวของแต่ละโรงพยาบาล โดยข้อมูลถูกนำมาใช้ในงานวิจัยประกอบไปด้วย ข้อความรีวิว วันที่รีวิว และคะแนนความพึงพอใจ ดังแสดงในรูปที่ ภาพประกอบ 7 หน้าเว็บไซต์การรีวิวของแต่ละโรงพยาบาล



ภาพประกอบ 6 ตัวอย่างข้อมูลรีวิวบนเว็บไซต์

คำอธิบายตัวอย่างข้อมูลรีวิว

หมายเลข 1 คือ ข้อมูลคะแนนความพึงพอใจของผู้ที่เคยใช้บริการ

หมายเลข 2 คือ ข้อมูลวันที่ของการแสดงความเห็น

หมายเลข 3 คือ ข้อมูลความคิดเห็นของผู้ที่เคยใช้บริการ

การเก็บรวบรวมข้อมูลจากเว็บไซต์โดยวิธีการสกัดข้อมูลด้วยเทคนิคการทำ web scraping ผ่านการใช้ไลบรารี Selenium และ BeautifulSoup ของ Python เพื่อทำการเปิด Browser ให้เข้าถึงแหล่งข้อมูลการรีวิวของทุกโรงพยาบาลบน www.honestdocs.com ด้วย chrome web drive โดยรายละเอียดการสกัดข้อมูลแบบอัตโนมัติ จะอธิบายเพิ่มเติมในส่วนของ ภาคผนวก

ข้อมูลที่ผ่านการสกัดมาจากเว็บไซต์จะถูกจัดเก็บอยู่ในรูปแบบตารางบนไฟล์ CSV ที่แบ่งข้อมูลตามแถวและคอลัมน์

ตาราง 4 รายละเอียดข้อมูลที่ได้รับจากการสกัดข้อมูล

ข้อมูล	คำอธิบาย
คะแนน	ข้อมูลคะแนนความพึงพอใจของผู้ใช้บริการโรงพยาบาล
เดือน/วัน/ปี เวลา	ข้อมูลวันที่ของการแสดงความเห็น
ข้อความรีวิว	ข้อมูลความคิดเห็นของผู้มีประสบการณ์การใช้บริการ
ชื่อโรงพยาบาล	ชื่อโรงพยาบาลที่ได้รับการรีวิว

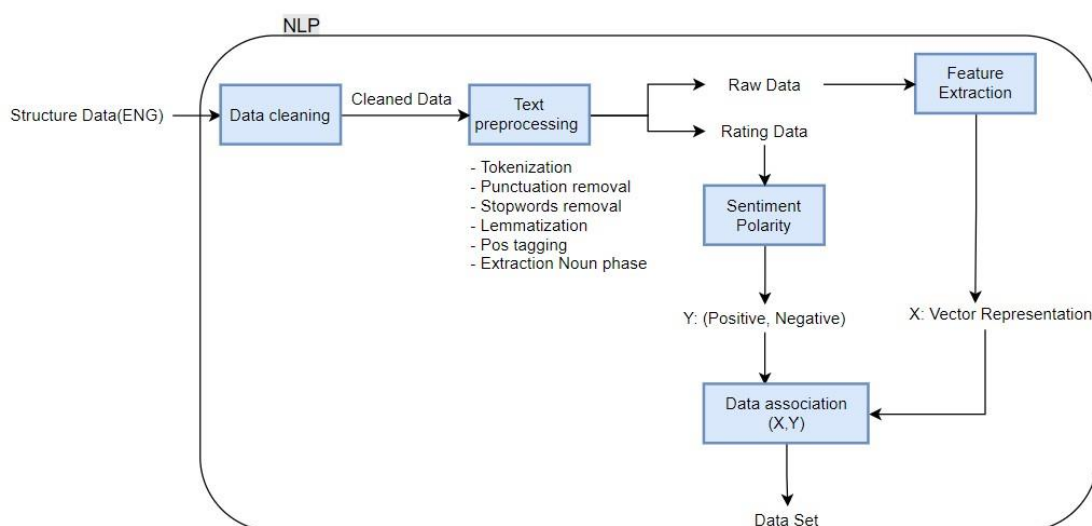
2.การแปลงข้อความรีวิวจากภาษาไทยเป็นภาษาอังกฤษคือ การนำไฟล์ที่ได้รับจากการสกัดข้อมูลมาอ่านด้วย ไลบรารี pandas ให้ข้อมูลจัดอยู่ในรูปแบบของ Data Frame และนำข้อมูลมาเข้าสู่กระบวนการลบสัญลักษณ์พิเศษ เช่น สติ๊กเกอร์อิโมจิ ออกจากข้อความรีวิวก่อนนำไปแปลเป็นภาษาอังกฤษด้วยไลบรารีของ Google translate และทำการจัดเก็บในรูปแบบของไฟล์ CSV

ตาราง 5 ตัวอย่างการแปลภาษาของข้อมูลรีวิว

Reviews_Th	Reviews_en
แพทย์ทุกคนรักษาอย่างดี	All doctors treated very well.

3.3 การประมวลผลภาษาธรรมชาติ

ในขั้นตอนของการเตรียมข้อมูล คือ การนำข้อมูลที่ผ่านการแปลภาษาแล้ว มาผ่านกระบวนการทำแปลงข้อมูลให้อยู่ในรูปแบบที่คอมพิวเตอร์สามารถประมวลได้และมีความพร้อมสำหรับการสร้างแบบจำลองการเรียนรู้ของเครื่องโดยประกอบไปด้วยขั้นตอนดังนี้



ภาพประกอบ 7 แผนผังการประมวลผลภาษาธรรมชาติ

1.การทำความสะอาดข้อความ (Text Cleaning) เป็นขั้นตอนที่สำคัญเพื่อการจัดการเตรียมข้อมูลประเภทข้อความให้อยู่ในรูปแบบที่ถูกต้องสำหรับการนำไปวิเคราะห์ และลดความซับซ้อนในการวิเคราะห์ ซึ่งช่วยเพิ่มประสิทธิภาพในการเรียนรู้ของเครื่องเพื่อสร้างแบบจำลองมีขั้นตอนดังนี้

2.การลบข้อความและอักขระพิเศษ (Remove none text and special character) จากการตรวจสอบข้อมูลพบว่ามีอักขระพิเศษ สัญลักษณ์ ผ่องอยู่ในข้อความ ซึ่งถือว่าไม่มีบริบทในข้อความ ไม่มีความหมาย และไม่ให้อข้อมูลที่ดีสำหรับแบบจำลองจึงจำเป็นต้องลบออก รวมถึง ลิงค์ และแฮชแท็กและอื่นๆควรจะถูกลบออกไปก่อนการเรียกใช้โมเดล

3.การแปลงตัวอักษรพิมพ์ใหญ่ให้เป็นตัวอักษรพิมพ์เล็กและการแปลงคำศัพท์ Plural เป็น Singular (Convert all text to lowercase and plural to singular) การแปลงตัวอักษรเพื่อหลีกเลี่ยงข้อผิดพลาดในระหว่างการฝึกฝนข้อมูลซึ่งอาจจะทำให้ได้รับการเรียนรู้ที่ต่างกันซึ่งตัวอักษรควรมีลักษณะเหมือนกันทั้งหมด

4.การตัดคำ (Word Segmentation) การแบ่งคำในประโยคในความคิดเห็นหนึ่งความคิดเห็นสามารถได้หลายประโยค และในประโยคหนึ่งประโยคประกอบไปด้วยคำหลายคำ แต่ไม่ใช่ทุกคำที่มีความหมาย ในการวิเคราะห์แต่ละคำจึงจำเป็นต้องแบ่งคำออกเป็นคำ ในข้อความภาษาอังกฤษจะแบ่งคำด้วย space bar สำหรับแต่ละประโยค

5.การลบคำหยุด (Remove Stop Word) เป็นคำที่เกิดขึ้นบ่อยๆในประโยคหรือในเอกสาร แต่ไม่มีความหมายและไม่มีประโยชน์เพื่อใช้ในการสร้างโมเดลเราจึงสามารถลบคำเหล่านี้ ออกรายการลิสต์ของคำศัพท์หรือกรองออกจากได้เลย เช่น a, an, the, also, just เป็นต้น

6.การแปลงคำศัพท์ให้อยู่ในพื้นฐานของคำศัพท์นั้น (Lemmatization) กระบวนการแปลงคำด้วยการอ้างอิงคำศัพท์ใน Dictionary ตามหลักไวยากรณ์ของภาษาใน Dictionary อย่างเหมาะสม เพื่อใช้ในการแปรผันคำเพื่อจำกัด inflection ของคำ เนื่องจากคำหนึ่งของในภาษาอังกฤษมีได้หลายรูปแบบที่มาจากรากศัพท์คำเดียวกัน

7.การติดแท็ก POS (Part-of-speech tagging) Parsing คือการตรวจสอบประโยคว่าถูกต้องตามกฎเกณฑ์หลักไวยากรณ์ของภาษาและการแสดง โครงสร้างของภาษา ด้วยโปรแกรม เรียกว่า Parser ซึ่งจะสร้างแผนภาพ เรียกว่า Parse Tree

3.4 การกำหนดประเภทความรู้สึก

การนำข้อมูลที่ผ่านมาการทำความสะอาดข้อมูลแล้ว มากำหนดประเภทความรู้สึกของแต่ละความคิดเห็น ซึ่งจะกำหนดโดยคะแนนความพึงพอใจของความคิดเห็นนั้น สามารถจัดประเภทความรู้สึกได้ 2 ประเภท คือ คะแนนความพึงพอใจ โดยคะแนนตั้งแต่ 1-2 จะกำหนดให้เป็นความรู้สึกประเภทเชิงลบ โดยแทนค่าในชุดข้อมูลด้วย 1 และคะแนนความพึงพอใจ 4-5 คะแนน จะกำหนดให้เป็นความรู้สึกประเภทเชิงบวก โดยแทนค่าในชุดข้อมูลด้วย 0

ตาราง 6 ตารางคลังข้อมูลรีวิวที่จัดประเภทความรู้สึก

คลังข้อมูล	จำนวน		
	ผลรวม	รีวิวเชิงบวก	รีวิวเชิงลบ
ข้อความรีวิว	4,987	1,056	3,931
ประโยค	15,718	4,501	11,217
คำ	120,681	34,377	86,304
คำที่ไม่ซ้ำ	8,300	3,306	4,994

3.5 การสร้างคุณลักษณะเพื่อการแปลงข้อมูล

ในงานวิจัยนี้ ในขั้นตอนการสร้างคุณลักษณะเพื่อแปลงข้อมูลสำหรับอัลกอริทึมที่เลือกใช้ 2 แบบ โดยตัวแปลงข้อมูลที่ใช้ในการประมวลผล คือ เทคนิค TFIDF สำหรับการเลือกใช้ในอัลกอริทึม Naïve Bayes และ Random Forest และใช้เทคนิค Word Embedding สำหรับการเลือกใช้อัลกอริทึมในการพัฒนาแบบจำลอง Long Short-Term Memory(LSTM)

3.5.1 การสร้างคุณลักษณะข้อมูลด้วยเทคนิค TFIDF

การสร้างฟีเจอร์สำหรับการเรียนรู้ของเครื่องด้วยอัลกอริทึม Multinomial Naïve Bayes และ Random forest จะใช้เทคนิคการแปลงข้อมูลด้วย การคำนวณ TF-IDF เพื่อวัดค่าทางสถิติที่ประเมินความเกี่ยวข้องของคำในข้อความรวิวกับทุกข้อความรีวิว

3.5.2 การสร้างคุณลักษณะข้อมูลด้วยเทคนิค word Embedding

ในงานวิจัยนี้ในการทดสอบสร้างแบบจำลองด้วยอัลกอริทึม LSTM มีการใช้เทคนิคการสร้างคุณลักษณะของข้อความด้วย word Embedding จากไลบรารี Keras ซึ่ง keras มีการใช้งานการฝึกฝนข้อมูลในการหาค่า ความสอดคล้อง คือ word2vec และจะได้การสร้างคุณลักษณะของคำศัพท์ ตามจำนวนคำที่เรากำหนด จากนั้นทำการแปลงข้อมูลให้อยู่ในรูปแบบคุณลักษณะที่คอมพิวเตอร์สามารถเข้าใจได้และนำข้อมูลเข้าสู่การแบ่งชุดข้อมูลเป็นชุดฝึกฝนและชุดทดสอบ

3.6 การแบ่งข้อมูลสำหรับชุดฝึกฝนและชุดทดสอบ

ข้อมูลที่ผ่านมากระบวนการแปลงข้อมูลหรือการสร้างฟีเจอร์ ข้อมูลจะถูกแบ่งออกเป็นสองส่วน เพื่อใช้ในการสร้างแบบจำลองการทำนาย และทดสอบประสิทธิภาพของแบบจำลองโดยแบ่งเป็นชุดข้อมูลฝึกฝน 80 เปอร์เซนต์ และชุดข้อมูลทดสอบ 20 เปอร์เซนต์

ตาราง 7 คลังข้อมูลที่แบ่งชุดข้อมูลฝึกฝนและชุดข้อมูลทดสอบ

คลังข้อมูล	ชุดข้อมูลฝึกฝน		ชุดข้อมูลทดสอบ	
	รีวิวเชิงบวก	รีวิวเชิงลบ	รีวิวเชิงบวก	รีวิวเชิงลบ
ประโยค	3,593	8,981	908	2,236
คำ	27,750	69,025	6,627	17,279
คำที่ไม่ซ้ำ	2,990	4,528	1,459	2,301

Positive class = Negative sentiment,

Negative class = Positive sentiment

คุณภาพชุดข้อมูลฝึกฝน (Data Quality)

เนื่องจากข้อมูลที่น่ามาใช้เป็นข้อมูลภายนอกและระบุประเภทของรีวิวด้วยข้อมูลการให้คะแนน ดังนั้นเรื่องความถูกต้องของข้อมูลเป็นสิ่งสำคัญสำหรับงานวิจัยนี้จึงรายงานความถูกต้องของข้อมูลที่น่ามาใช้เป็นข้อมูลชุดฝึกฝนโดยมีการตรวจสอบโดย สุ่มข้อมูลรีวิว 100 รีวิว เนื่องจากข้อมูล 1 รีวิว สามารถมีได้หลายประโยคและประโยคนั้นจะถูกระบุประเภทเดียวกันทั้งหมดของ 1 รีวิว เพื่อตรวจสอบข้อมูลว่ามีความถูกต้องในการติดฉลากในชุดข้อมูลฝึกฝนมากแค่ไหน เพื่อให้ทราบถึงความสอดคล้องระหว่างค่าความถูกต้องของข้อมูลและการวัดประสิทธิภาพของแบบจำลอง ซึ่งจะส่งผลต่อค่าความถูกต้องของโมเดล จึงแสดงคุณภาพของข้อมูลได้ดังนี้

ตาราง 8 ข้อมูลแสดงคุณภาพของข้อมูลจากแหล่งข้อมูล

คุณภาพของข้อมูล (Data Quality)	
ข้อมูลที่ดีฉลากถูกต้อง	ข้อมูลที่ดีฉลากไม่ถูกต้อง
57%	43%

3.7 การสร้างแบบจำลองการแยกประเภทความรู้สึก

การสร้างแบบจำลองการเรียนรู้ของเครื่องเพื่อใช้ในการแยกประเภทความรู้สึกในงานวิจัยนี้อัลกอริทึมที่เราสนใจนำมาเปรียบเทียบประสิทธิภาพของแบบจำลองกับชุดข้อมูลรีวิวจากประสบการณ์การให้บริการของผู้ป่วย ด้วยอัลกอริทึม Naïve Bayes, Random forest และ LSTM

3.7.1 การสร้างแบบจำลองด้วยอัลกอริทึม Random Forest

ในงานวิจัยนี้ได้มีการใช้วิธีการการแยกประเภทด้วยการเรียนรู้ของอัลกอริทึมที่มีการสร้างชุดชุดตัวตัดสินใจ และมีการสุ่มจัดชุดข้อมูลกระจายไปยังแต่ละชุดตัวตัดสินใจ และแยกการทำนายจากนั้นโหวตจากชุดการตัดสินใจหลายๆตัว ซึ่ง random forest เป็นอัลกอริทึมการเรียนรู้แบบ Ensemble learning เป็นอัลกอริทึมที่เป็นประเภทการเรียนรู้ที่มีอัลกอริทึมเดียวกันหลายๆชุดหรือมีการตัดสินใจหลายๆครั้งเพื่อเพิ่มความถูกต้องในการทำนายของแบบจำลองโดยอัลกอริทึมนี้สามารถใช้ในงาน การแยกประเภทและการทำนายค่าความถดถอย รูปแบบการทำงานของ RF สามารถอธิบายได้ว่า เป็นกลุ่มของตัวแยกประเภทที่เป็นโครงสร้างแบบต้นไม้การตัดสินใจ ด้วยการโหวตในรูปแบบของ bagging โดยจะใช้การเลือกผลโหวตจากผลลัพธ์ที่มีการทำนายตรงกันมากที่สุด ซึ่งกลยุทธ์นี้จะช่วยให้ RF มีค่าความแม่นยำที่ดีและใช้เวลาเร็วในการสร้างแบบจำลองในงานวิจัยนี้ได้ใช้การสร้างแบบจำลองโดยใช้พารามิเตอร์ ดังนี้



ตาราง 9 ตารางแสดงค่าพารามิเตอร์ของการสร้างแบบจำลองด้วยอัลกอริทึม Random Forest

ชื่อพารามิเตอร์	ค่าพารามิเตอร์
n_estimators	150
criterion	True
max_depth	8
min_samples_split	2
min_samples_leaf	1
min_weight_fraction_leaf	0.0
max_features	auto
max_leaf_nodes	None
min_impurity_decrease	0.0
min_impurity_split	None
bootstrap	True
oob_score	False
n_jobs	None
random_state	None
verbose	0
warm_start	False
class_weight	None
ccp_alpha	0.0
max_samples	None

3.7.2 การสร้างแบบจำลองด้วยอัลกอริทึม Multinomial Naive Bayes

อัลกอริทึม Multinomial Naive Bayes เป็นวิธีการเรียนรู้ความน่าจะเป็นที่ส่วนใหญ่ใช้ในการประมวลผลภาษาธรรมชาติ (NLP) อัลกอริทึมนี้ใช้ทฤษฎี Bayes และคาดการณ์แท็กของข้อความ เช่น ข้อความบนอีเมลหรือบทความในหนังสือพิมพ์ จะคำนวณความน่าจะเป็นของแต่ละแท็กสำหรับตัวอย่างที่กำหนด จากนั้นให้แท็กที่มีความน่าจะเป็นสูงสุดเป็นผลลัพธ์

ตัวแยกประเภท Naive Bayes คือชุดของอัลกอริทึมจำนวนมากที่อัลกอริทึมทั้งหมดมีหลักการร่วมกันเพียงข้อเดียว และการมีอยู่ของคุณสมบัติเฉพาะในคลาสนั้นไม่เกี่ยวข้องกับการมีอยู่ของคุณสมบัติอื่นใด โดยพารามิเตอร์ที่ใช้ในการสร้างแบบจำลองประกอบคือ

ตาราง 10 ตารางแสดงค่าพารามิเตอร์ของการสร้างแบบจำลองด้วยอัลกอริทึม Naive Bayes

ชื่อพารามิเตอร์	ค่าพารามิเตอร์
Alpha	1
fit_prior	True
class_prior	none

3.7.3 การสร้างแบบจำลองการวิเคราะห์ด้วย Long Short-Term memory

อัลกอริทึม Long short-term memory (LSTM) เป็นเครือข่ายโครงสร้างประสาทเทียมที่ต่อยอดมาจาก recurrent neural network (RNN) แต่จะดีกว่า RNN ไม่มีปัญหาในการอัปเดตน้ำหนัก เมื่อทำงานกับข้อมูลที่เมื่อดำเนินยาวๆ ปกติการทำงานของ RNN เป็นเหมือน Neural Network ที่มี memory ธรรมดาอยู่ข้างในเพื่อบันทึก Hidden state ซึ่ง LSTM ก็มี memory ที่เหนือกว่า RNN คือมีตัวบอกได้ว่า สามารถ write forget read เมื่อไร

ตาราง 11 ตารางแสดงค่าพารามิเตอร์ของการสร้างแบบจำลองด้วยอัลกอริทึม LSTM

ชื่อพารามิเตอร์	ค่าพารามิเตอร์
Max feature	2000
Optimizer	adam
จำนวน LSTM Layer	4
Learning rate	0.01
Batch size	10

ตาราง 12 ตารางผลลัพธ์ Summary ของโมเดล LSTM ที่ได้จากการใช้พารามิเตอร์ข้างต้น

Layer (Type)	Output Shape	Param#
Embedding	(none, 250, 128)	256000
Spatial_ dropout1d	(none, 250,128)	0
LSTM	(none, 196)	124800
Dense	(none,2)	394
Total params: 511,194		
Trainable params: 511,194		
Non-trainable params: 0		

แบบจำลองการวิเคราะห์ความรู้สึกด้วยคุณลักษณะ Word Embedding และอัลกอริทึมอัลกอริทึม Word Embedding ทำการสร้างคุณลักษณะของข้อความขึ้นมาตามจำนวนคำศัพท์ที่เราได้ทำการกำหนดขึ้นและจำนวนมิติของ Word Embedding เมื่อได้คุณลักษณะของข้อมูลแล้วจะแปลงข้อมูลให้อยู่ในรูปของคุณลักษณะที่สามารถนำมาสร้างแบบจำลองการทำนายผลการวิเคราะห์ความรู้สึกด้วยอัลกอริทึม Long Short Term Memory จากนั้นทำการทดสอบประสิทธิภาพการทำนายและบันทึกผลบันทึกผล

3.8 การประเมินผลประสิทธิภาพแบบจำลองการเรียนรู้ของเครื่อง

ในการวัดประสิทธิภาพของแบบจำลองในการทำนายความรู้สึกด้วยการแสดงผลการรายงานค่าทาง เนื่องจากข้อมูลรีวิวประสบการณ์การใช้บริการโรงพยาบาลที่ให้ความหมายเชิงด้านลบ (Negative Sentiment) ซึ่งมีความสำคัญต่อระบบการจัดการบริหารคุณภาพโรงพยาบาลมากกว่า ข้อมูลรีวิวที่ให้ความหมายในเชิงบวก (Positive Sentiment) เราจึงให้ความสำคัญกับข้อมูล การรีวิวในเชิงลบ โดยข้อมูลเหล่านี้จะทำให้เราทราบถึงปัญหาที่เกิดขึ้น หรือคำติชมในระบบจัดการบริหารในด้านต่างๆ ได้อย่างรวดเร็ว เพื่อตอบสนองวัตถุประสงค์ของผู้กำหนดนโยบายของโรงพยาบาล จึงเลือกการประเมินประสิทธิภาพโมเดลด้วยค่าทางสถิติในรูปแบบของค่า Recall เพื่อให้เหมาะสมและตรงตามวัตถุประสงค์ของงานวิจัย การแสดงผลของค่าทางสถิติ ประกอบไปด้วยค่า Accuracy, Precision, Recall และ F1- Score และยังประเมินในด้านของระยะเวลาที่ใช้ในการสร้างแบบจำลองด้วย ซึ่งในการคำนวณค่าทางสถิติของตัวชี้วัดสามารถแสดงสมการได้ดังนี้

Accuracy คือ เป็นการวัดความถูกต้องของแบบจำลอง โดยพิจารณารวมทุกคลาส จากการเปรียบเทียบข้อมูลในการทำนายระหว่างผลลัพธ์จากการทำนาย กับผลลัพธ์เป้าหมายว่ามีความสัมพันธ์กันอย่างไร ซึ่งดูจากการสร้างเมทริกซ์ประสิทธิภาพ

$$Accuracy = \frac{(TPs + TNS)}{(TPs + TNS + FPs + FNs)} \quad (8)$$

Confusion Matrix คือตารางสำคัญในการวัดประสิทธิภาพของแบบจำลองการเรียนรู้ของเครื่องในการแก้ปัญหาการแยกประเภท

True Positive คือ แบบจำลองสามารถทำนายได้ว่า “จริง” และมีค่าเป็น “จริง”

True Negative คือ แบบจำลองสามารถทำนายได้ว่า “ไม่จริง” และมีค่าเป็น “ไม่จริง”

False Positive คือ แบบจำลองสามารถทำนายได้ว่า “จริง” และมีค่าเป็น “ไม่จริง”

False Negative คือ แบบจำลองสามารถทำนายได้ว่า “ไม่จริง” และมีค่าเป็น “จริง”

Precision คือ เป็นการวัดความแม่นยำของข้อมูล โดยการคำนวณจาก TP เทียบกับกลุ่มของข้อมูลที่ทำนายเป็นด้านบวกทั้งหมด (TP + FP) จากสูตรสมการดังนี้

$$Precision = \frac{TPs}{(TPs + FPs)} \quad (9)$$

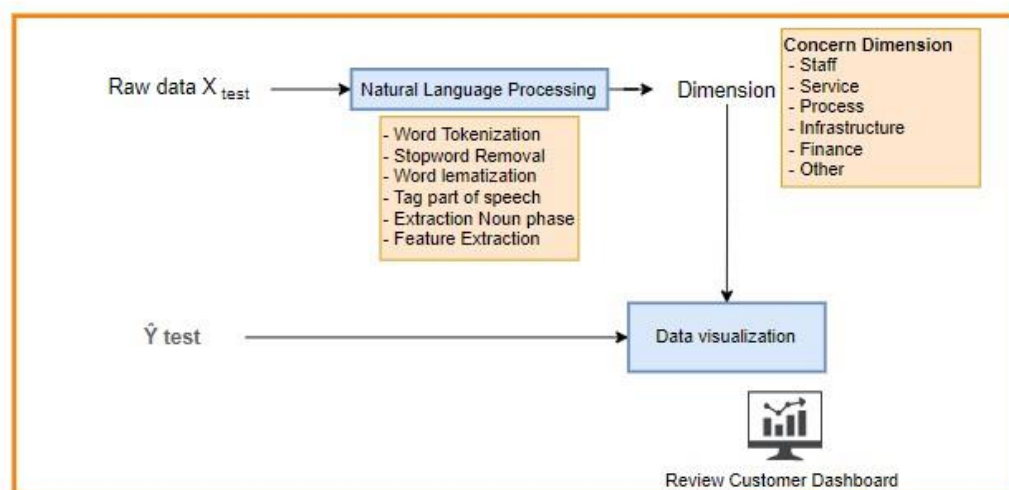
Recall คือ ค่าความไว หรือค่าระลึก เป็นการวัดความถูกต้องของแบบจำลอง โดยคำนวณจากค่า TP เทียบกับกลุ่มของข้อมูลที่มีผลลัพธ์เป็นจริงทั้งหมด จากสูตรสมการดังนี้

$$Recall = \frac{TPs}{TPs + FNs} \quad (10)$$

F1-Score คือ (ค่าประสิทธิภาพโดยรวม) นำสองค่า (Precision และ Recall) มาพิจารณาร่วมกันเพื่อความแม่นยำ จากสูตรสมการดังนี้

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (11)$$

3.9 การประยุกต์ใช้แบบจำลองการเรียนรู้เพื่อแยกประเภทความรู้สึกและหาข้อมูลเชิงลึก



ภาพประกอบ 8 แผนผังกระบวนการทำ Sentiment Analysis

ในการศึกษาการวิเคราะห์ความรู้สึกของข้อความรีวิวจากประสบการณ์การใช้บริการของโรงพยาบาลรัฐและโรงพยาบาลเอกชนมีจุดประสงค์เพื่อต้องการสร้างรายงานการแสดงผลการวิเคราะห์ข้อมูลของผู้ป่วยที่เข้ามาใช้บริการในโรงพยาบาลเพื่อประเมินคุณภาพของระบบการจัดการบริหารโรงพยาบาลในด้านต่างๆตามหัวข้อตัวชี้วัดที่เราสนใจและทำการเปรียบเทียบข้อมูลระหว่างโรงพยาบาลรัฐและเอกชนเพื่อหาข้อมูลเชิงลึกและรับรู้ถึงปัญหาที่เกิดขึ้นกับคุณภาพการจัดการในโรงพยาบาลทั้งภาครัฐและเอกชน ซึ่งมีกระบวนการดำเนินงานวิเคราะห์ความรู้สึกข้อความรีวิวของผู้ป่วยในโรงพยาบาลมี ดังนี้

ขั้นตอนที่ 1 ข้อมูลรีวิวจากประสบการณ์การใช้บริการของโรงพยาบาลรัฐและเอกชน

ขั้นตอน 2 การแยกประเภทความรู้สึกด้วยแบบจำลองการเรียนรู้ของเครื่อง

ขั้นตอนที่ 3 การจำแนกหัวข้อตัวชี้วัดของข้อความรีวิว

ขั้นตอนที่ 4 การนำเสนอแผนภาพข้อมูล

3.9.1 ข้อมูลรีวิวจากประสบการณ์การใช้บริการของโรงพยาบาลรัฐและเอกชน

ข้อมูลรีวิวประสบการณ์การใช้บริการโรงพยาบาลที่นำมาใช้ในการวิเคราะห์เพื่อสร้างการรายงานผลเป็นข้อมูลการรีวิวของโรงพยาบาลขนาดใหญ่และมีชื่อเสียงใน 3 อันดับแรกของประเทศของโรงพยาบาลเอกชนและโรงพยาบาลรัฐในประเทศไทย

ตาราง 13 ข้อมูลรีวิวสำหรับการวิเคราะห์ความรู้สึกจากประสบการณ์การใช้บริการโรงพยาบาลในประเทศไทย

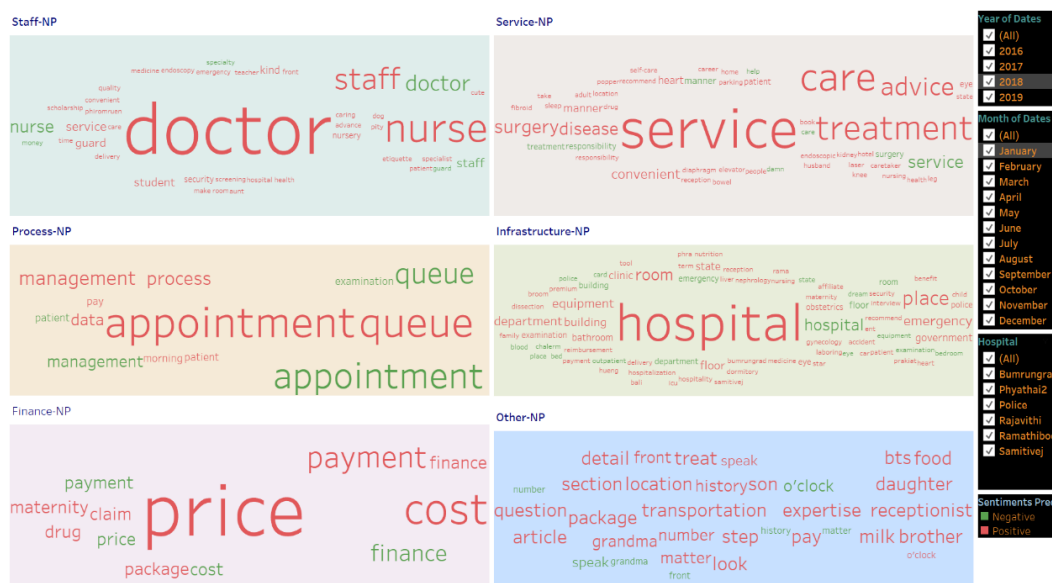
ข้อมูลรีวิวโรงพยาบาล	จำนวนข้อมูล
โรงพยาบาลรัฐ 1	80
โรงพยาบาลรัฐ 2	590
โรงพยาบาลรัฐ 3	731
โรงพยาบาลเอกชน 1	583
โรงพยาบาลเอกชน 2	95
โรงพยาบาลเอกชน 3	347

3.9.2 การแยกประเภทความรู้สึกด้วยแบบจำลองการเรียนรู้ของเครื่อง

ความรู้สึกของข้อมูลรีวิวการเตรียมข้อมูลเพื่อนำเข้าสู่แบบจำลองจะผ่านกระบวนการใช้เทคนิคการประมวลผลภาษาธรรมชาติเพื่อให้คอมพิวเตอร์สามารถเข้าใจและเรียนรู้ได้ และเนื่องจากการทดสอบและเลือกแบบจำลองที่มีประสิทธิภาพสูงสุดในการทำนายข้อมูลประสบการณ์การใช้บริการโรงพยาบาล จึงได้นำมาประยุกต์ใช้ในการวิเคราะห์ความรู้สึกของชุดข้อมูลใหม่ของการรีวิวประสบการณ์ของผู้มาใช้บริการของโรงพยาบาลรัฐและโรงพยาบาลเอกชน ที่ต้องการวิเคราะห์ข้อมูลเพื่อนำไปใช้ในการหาข้อมูลเชิงลึก จากการเตรียมข้อมูล

3.9.3 การจำแนกหัวข้อตัวชี้วัดของข้อความรีวิว

การจำแนกหัวข้อตัวชี้วัด (Concern Dimension) สำหรับระบบการจัดการ คุณภาพของโรงพยาบาลที่สนใจ ที่กำหนดขึ้นเพื่อใช้ในการระบุข้อความรีวิวที่กล่าวถึงคุณภาพด้านต่างๆของโรงพยาบาล ในงานวิจัยนี้กำหนดไว้ 6 ด้าน คือ ด้านการบริการ ด้านโครงสร้างโรงพยาบาล ด้านการดำเนินงาน ด้านการเงิน ด้านเจ้าหน้าที่ และอื่นๆ ซึ่งสามารถระบุได้โดยการค้นหาคำที่ปรากฏ



ภาพประกอบ 10 แดชบอร์ดการแสดงผล word cloud ในหัวข้อตัวชี้วัดแต่ละด้าน

ข้อมูลที่นำมาใช้ในการสร้างแดชบอร์ดเพื่อแสดงรายงานผลการวิเคราะห์ คือข้อมูลรีวิวกี่ ระบุความรู้สึกในเชิงบวกและเชิงลบของผู้ใช้บริการผ่านการใช้เทคนิคการเรียนรู้ของเครื่องและมีการจัดหมวดหมู่ของคำที่ปรากฏในข้อความรีวิวให้ตรงกับหัวข้อปัญหาแล้ว จากนั้น เราใช้เครื่องมือ Tableau เพื่อการสร้างแดชบอร์ดและแสดงรายงานผลการวิเคราะห์ข้อมูลที่สำคัญโดยมีแนวคิดและการสร้างแดชบอร์ดที่เหมาะสมดังนี้

Box Text คือ การจัดตำแหน่งไว้อยู่ด้านบนซ้ายของแดชบอร์ด เพื่อแสดงจำนวนข้อความรีวิวทั้งหมด จำนวนข้อความรีวิวในเชิงบวกและจำนวนข้อความรีวิวในเชิงลบ

Trend Chart คือ การใช้เป็นกราฟแสดงแนวโน้มของความถี่ของการรีวิวในการพูดถึงโรงพยาบาลในเชิงลบและเชิงบวกมากขึ้นหรือน้อยแค่ไหน ซึ่งมีประโยชน์กับผู้กำหนดนโยบายของโรงพยาบาลที่จะสามารถเห็นถึงปัญหาที่ผู้ใช้บริการได้รับและดำเนินการแก้ปัญหาได้ถูกจุดและเหมาะสมกับช่วงเวลา ซึ่งเราสามารถเลือกช่วงเวลาดูแนวโน้มได้ เป็นรายวัน รายเดือน รายไตรมาส และรายปีได้

Word Cloud คือ การใช้ word cloud เพื่อแสดงความถี่ของคำที่ปรากฏในหัวข้อตัวชี้วัดที่สนใจจากข้อความรีวิว เพื่อดูว่าผู้ใช้บริการส่วนใหญ่กล่าวถึงคำไหนมากที่สุดในการรีวิว

Pie Chart คือ การใช้ pie chart แสดงสัดส่วนของจำนวนรีวิวที่ถูกพูดถึงในรูปแบบสัดส่วนที่เป็นเปอร์เซ็นต์ตามหัวข้อตัวชี้วัดที่สนใจที่แสดงในข้อมูลรีวิวเชิงบวกและข้อมูลรีวิวเชิงลบ

Filter คือ การใช้ฟิลเตอร์เพื่อแสดงผลการรายงานในแดชบอร์ดนี้มี 4 ส่วน คือ 1. ฟิลเตอร์เลือกช่วงวันเวลาเป็น รายปี รายเดือน ฟิลเตอร์นี้ทำให้เราสามารถเลือกช่วงระยะเวลาเพื่อเรียกดูข้อมูลผลการวิเคราะห์ได้ตามที่ต้องการ 2. ฟิลเตอร์เลือกโรงพยาบาล เราสามารถเลือกโรงพยาบาลที่เราสนใจเพื่อดูรายงานผลการวิเคราะห์และสามารถนำข้อมูลของแต่ละโรงพยาบาลมาเปรียบเทียบกัน 3. ฟิลเตอร์ของหัวข้อตัวชี้วัดที่สนใจตามระบบการจัดการคุณภาพของโรงพยาบาล โดยทั้ง 6 หัวข้อนี้เราสามารถเลือกหัวข้อตัวชี้วัดที่สนใจเพื่อดูรายละเอียดข้อมูลในส่วนอื่นของแดชบอร์ดได้ 4. ฟิลเตอร์ประเภทความรู้สึก ใช้เพื่อเลือกดูข้อมูลที่กำลังถึงข้อมูลรีวิวเชิงบวกหรือข้อมูลรีวิวเชิงลบ ฟิลเตอร์ทั้งหมดนี้จะสามารถเชื่อมถึงกันได้ทุกกราฟบนแดชบอร์ดซึ่งเรียกว่า Interactive Dashboard



บทที่ 4

ผลการดำเนินการวิจัย

ความสำคัญของงานวิจัยนี้ในการวิเคราะห์ความรู้สึกข้อมูลประเภทข้อความของข้อมูล รีวิวจากประสบการณ์การใช้บริการของผู้ป่วยในโรงพยาบาลด้วยการสร้างแบบจำลองการเรียนรู้ของเครื่อง และการสร้างรายงานการแสดงผลเพื่อหาข้อมูลเชิงลึก โดยผลของงานวิจัยจะแบ่งเป็นสองส่วนที่ประกอบไปด้วย ผลการประเมินประสิทธิภาพจากการเปรียบเทียบแบบจำลองทั้งสาม อัลกอริทึม และผลการวิเคราะห์ข้อมูลเชิงลึกจากรายงานผลบนแดชบอร์ดจากการประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่องในการจำแนกประเภทความรู้สึก

การเปรียบเทียบผลการประเมินประสิทธิภาพของแบบจำลองที่นำมาทดสอบทั้งสาม อัลกอริทึมได้แก่ Multinomial Naive Bayes, Random forest และ LSTM พบว่าแบบจำลองที่สามารถทำนายข้อความรีวิวได้มีความถูกต้องที่สุดด้วยเมตริกซ์ที่สำคัญสำหรับการวัดประสิทธิภาพของแบบจำลองด้วยค่า Recall โดยเรียงจากมากไปน้อย คือ Random forest Multinomial Naïve Bayes และ LSTM ที่ 0.71, 0.62 และ 0.46 ตามลำดับ จากผลการประเมิน Random Forest ให้ประสิทธิภาพสูงสุดในการจำแนกประเภทความรู้สึก ซึ่งสามารถอธิบายได้ว่า โมเดล Random Forest แบบจำลองเป็นประเภทการเรียนรู้แบบ Ensemble ที่การทำงานของแบบจำลองสร้าง decision Tree หลายๆโมเดล ซึ่งผลลัพธ์ที่ได้จะได้รับการคำนวณผลการทำนายจากการโหวตผลลัพธ์ที่ถูกเลือกจาก decision Tree มากที่สุด และเกิดจากการทำงานร่วมกันของหลายๆโมเดล จึงช่วยให้การจำแนกประเภทความรู้สึกได้ความแม่นยำมากกว่าการจำแนกแบบโมเดลเดียว โมเดล Multinomial Naïve Bayes (NB) เป็นโมเดลที่นิยมใช้เป็น Baseline ในงานวิเคราะห์ข้อความ โดยคุณสมบัติของ NB ประสิทธิภาพของโมเดลไม่ได้ขึ้นอยู่กับปริมาณของข้อมูล และเนื่องจากการเตรียมข้อมูลเป็น Unique Word จากการสร้างคุณลักษณะข้อมูลด้วย TFIDF ด้วย Bag of word ทำให้ข้อมูลเป็นลักษณะ cadential Independent ซึ่งแต่ละค่าจะไม่เป็น Dependent กัน จึงอาจจะทำให้ประสิทธิภาพยังไม่ดีเท่าที่ควร เป็นแบบจำลองที่ใช้เป็น Baseline ในงานการวิเคราะห์ความรู้สึกกับข้อมูลประเภทข้อความ และเนื่องจากข้อมูลที่ใช้ในการสร้างโมเดลใช้วิธีการสร้างคุณลักษณะแบบ Bag of word โดยสร้างด้วยการใช้เทคนิค TFIDF ซึ่งข้อมูลจะไม่มีความเป็นลำดับ (Independently) และการวิเคราะห์ข้อมูลจะไม่ขึ้นแก่กัน จึงทำให้ผลที่ได้ยังไม่ดีพอ โมเดล LSTM เป็นโมเดลที่ต้องการปริมาณข้อมูลมากเมื่อทำการตรวจสอบจากจำนวนพารามิเตอร์ที่ใช้แล้วมีปริมาณเยอะกว่าข้อมูลนำเข้ามากๆ โดยปกติการ

ทำงานของ LSTM ควรจะมีข้อมูลจำนวนมากกว่า parameter สำหรับการฝึกฝนโมเดล จึงทำให้ประสิทธิภาพในการทำงานของโมเดลไม่ดีเท่าที่ควร

ตาราง 15 ตารางรายงานผลค่าทางสถิติ Accuracy, Precision, Recall และ F1-Score.

อัลกอริทึม	ค่าความถูกต้อง	ค่าความแม่นยำ	ค่าความระลึก	ค่าเฉลี่ย
Naive Bayes	0.77	0.79	0.62	0.64
Random Forest	0.65	0.68	0.71	0.65
LSTM	0.47	0.46	0.46	0.44

ผลจากการประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่องเพื่อทำนายความรู้สึกจากการรีวิวของผู้ใช้บริการเพื่อนำมาใช้ประโยชน์ในการทำระบบเบื้องหลังของการสร้างแดชบอร์ดเพื่อการแสดงผลรายงานการวิเคราะห์ข้อมูลรีวิวเพื่อศึกษาการหาข้อมูลเชิงลึกของข้อมูลการรีวิวเพื่อนำมาประเมินระบบการจัดการคุณภาพของโรงพยาบาลรัฐและโรงพยาบาลเอกชนผ่านการรายงานการแสดงผลแดชบอร์ด สามารถสรุปข้อมูลที่ได้จากรายงานผลการวิเคราะห์จากการเก็บรวบรวมข้อมูลรีวิวตั้งแต่ปี 2016 -2019 ได้ดังนี้

4.1 การวิเคราะห์จำนวนข้อมูลรีวิวเชิงลบและเชิงบวกบนกราฟแท่ง



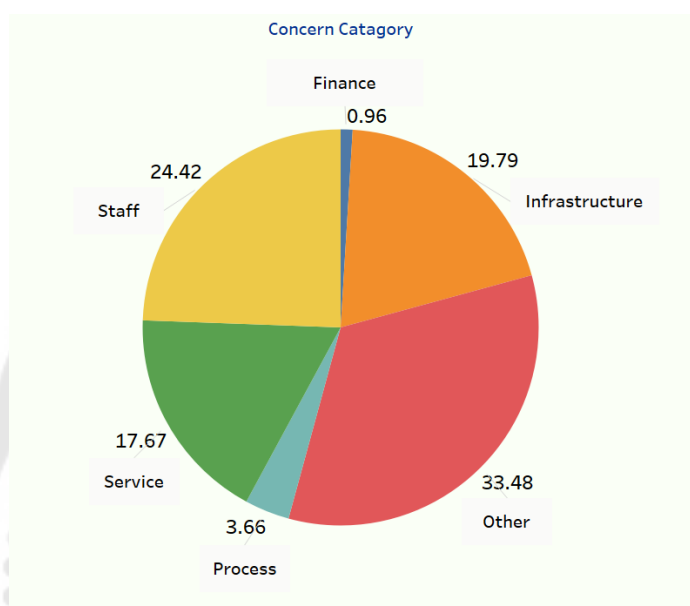
ภาพประกอบ 13 กราฟแท่งแสดงสัดส่วนรีวิวเชิงลบและเชิงบวกของโรงพยาบาลเอกชน



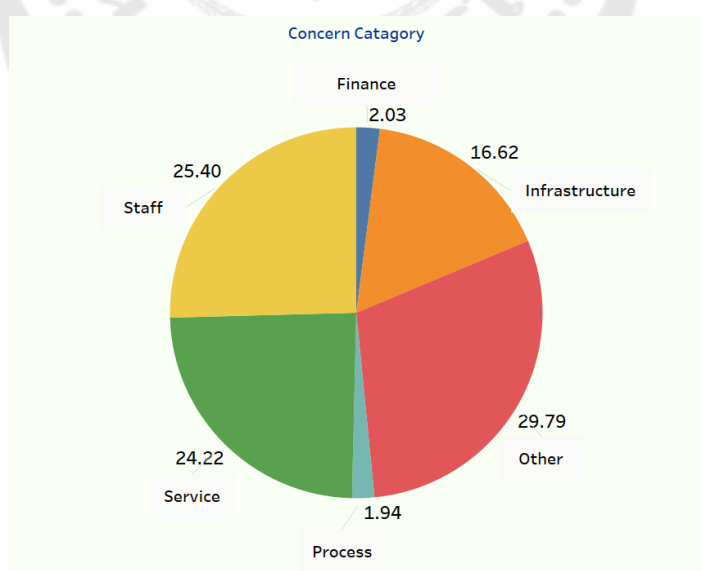
ภาพประกอบ 14 กราฟแท่งแสดงสัดส่วนรีวิวเชิงลบและเชิงบวกของโรงพยาบาลรัฐ

กราฟแท่งแสดงจำนวนรวิวในเชิงลบและเชิงบวกเปรียบเทียบระหว่างโรงพยาบาลรัฐและเอกชนจะเห็นว่าสัดส่วนข้อมูลเชิงลบของโรงพยาบาลรัฐมีมากกว่าโรงพยาบาลเอกชน

4.2 การวิเคราะห์ข้อมูลสัดส่วนของหัวข้อตัวชี้วัดบนกราฟพาย



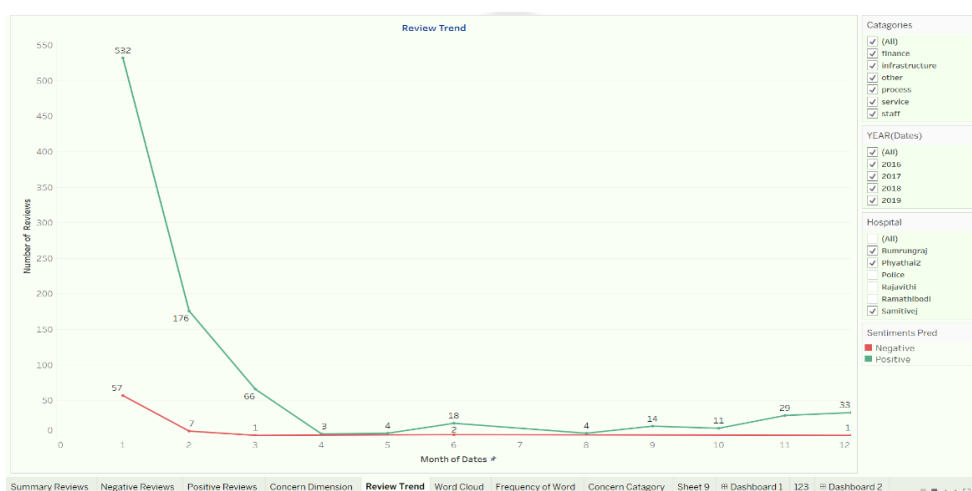
ภาพประกอบ 15 กราฟพายแสดงสัดส่วนรวิวของโรงพยาบาลรัฐตามหัวข้อตัวชี้วัด



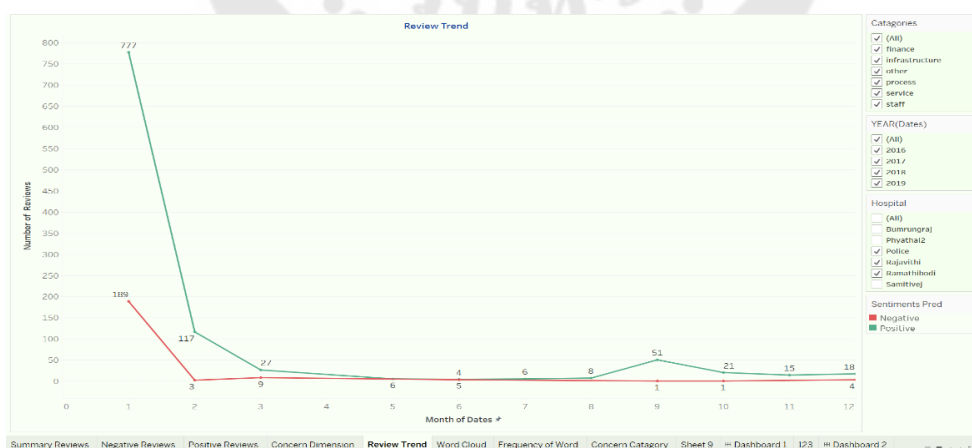
ภาพประกอบ 16 กราฟพายแสดงสัดส่วนรวิวของโรงพยาบาลเอกชนตามหัวข้อตัวชี้วัด

ผลการวิเคราะห์ข้อมูลจากแดชบอร์ดบนกราฟพายแบ่งสัดส่วนออกเป็นทั้งหมด 6 หัวข้อ ตัวชี้วัด พบว่าข้อมูลรีวิวกมีการกล่าวถึงในประเด็นของหัวข้อตัวชี้วัดด้านอื่นๆ มีมากที่สุดทั้งในโรงพยาบาลเอกชน และ โรงพยาบาลรัฐ แต่เมื่อเทียบในด้านบริการ จะเห็นว่าโรงพยาบาลเอกชนมีการพูดถึงมากกว่า โรงพยาบาลรัฐ ส่วนในด้านของการเงิน ด้านการดำเนินการ ด้านบุคลากร ด้านโครงสร้าง มีสัดส่วนในการพูดถึงใกล้เคียงกันทั้ง โรงพยาบาลเอกชน และ โรงพยาบาลรัฐ

4.3 การวิเคราะห์จำนวนข้อมูลรีวิวเชิงบวกและเชิงลบในช่วงเวลารายเดือนบนกราฟเส้น



ภาพประกอบ 17 กราฟเส้นแสดงจำนวนรีวิวเชิงบวกและลบของโรงพยาบาลเอกชน



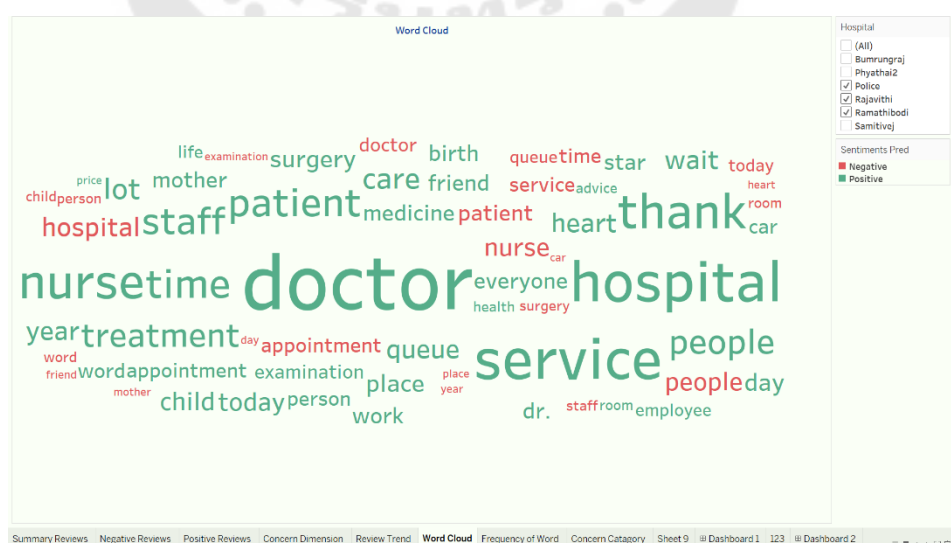
ภาพประกอบ 18 กราฟเส้นแสดงจำนวนรีวิวเชิงบวกและลบของโรงพยาบาลรัฐบาล

ผลการวิเคราะห์ข้อมูลจากแดชบอร์ดบนกราฟเส้นแบบรายเดือนของข้อมูลรีวิวเชิงลบ และเชิงบวก พบว่า ในทั้งสอง รูปแบบโรงพยาบาลมีการรีวิวเชิงบวกมากกว่าเชิงลบ และการรีวิวจะเกิดความถี่มากในช่วงต้นปี และมีแนวโน้มลดลงตั้งแต่กลางปีจนถึงท้ายปีจึงมีการเพิ่มขึ้นในช่วงท้ายปี

4.4. การวิเคราะห์คำที่ปรากฏในข้อความรีวิวแสดงผ่านเวิร์ดคลาวด์



ภาพประกอบ 19 เวิร์ดคลาวด์แสดงความถี่คำที่ปรากฏในข้อความรีวิวของโรงพยาบาลเอกชน



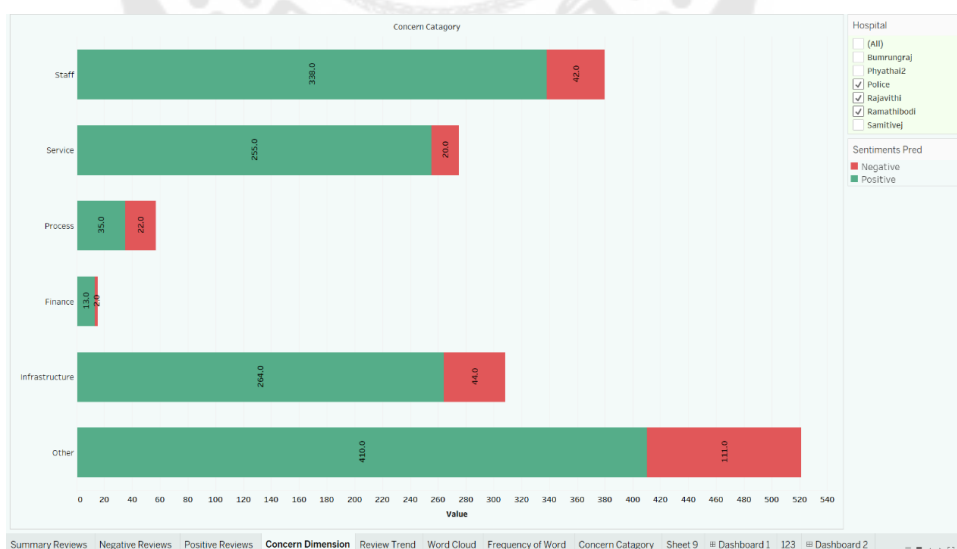
ภาพประกอบ 20 เวิร์ดคลาวด์แสดงความถี่คำที่ปรากฏในข้อความรีวิวของโรงพยาบาลรัฐบาล

ผลการวิเคราะห์ข้อมูลจากแดชบอร์ดบนเวิร์ดคลาวด์จะพบว่าข้อความส่วนมากที่ปรากฏจะเป็นข้อความเชิงบวกแต่ถ้าเมื่อเปรียบเทียบกันระหว่าง โรงพยาบาลเอกชน และ โรงพยาบาลรัฐ จะพบว่าคำที่ปรากฏของโรงพยาบาลรัฐ จะมีปรากฏคำที่อยู่ในรีวิวเชิงลบมากกว่าโรงพยาบาลเอกชน

4.6 การวิเคราะห์จำนวนข้อมูลรีวิวเชิงลบและเชิงบวกในแต่ละหัวข้อตัวชี้วัด



ภาพประกอบ 21 กราฟแท่งแสดงจำนวนข้อมูลรีวิวเชิงลบและเชิงบวกของโรงพยาบาลเอกชน



ภาพประกอบ 22 กราฟแท่งแสดงจำนวนข้อมูลรีวิวเชิงลบและเชิงบวกของโรงพยาบาลรัฐ

ผลการวิเคราะห์ข้อมูลจากแดชบอร์ดบนกราฟแท่งจะพบว่าข้อมูลในแต่ละด้านในเชิงบวกและเชิงลบ พบว่ารวิวิของโรงพยาบาลเอกชนมีอัตราส่วนของรวิวิเชิงลบน้อยกว่ารวิวิของโรงพยาบาลรัฐ ในทุกๆหัวข้อตัวชี้วัด

ตาราง 16 ข้อมูลจำนวนรวิวิเชิงลบและรวิวิเชิงบวกในหัวข้อตัวชี้วัดของโรงพยาบาลรัฐและเอกชน

หัวข้อตัวชี้วัด	จำนวนรวิวิ	
	โรงพยาบาลรัฐ	โรงพยาบาลเอกชน
ด้านเจ้าหน้าที่	42	13
ด้านการบริการ	20	8
ด้านการดำเนินงาน	22	8
ด้านการเงิน	2	1
ด้านโครงสร้าง	44	7
ด้านอื่นๆ	111	39

จากการวิเคราะห์ข้อมูลรวิวิเชิงลบในหัวข้อตัวชี้วัด พบว่าโรงพยาบาลเอกชนพบปัญหามากที่สุดคือในด้าน เจ้าหน้าที่ (staff) การดำเนินงาน (Process) การบริการ (service) โครงสร้างโรงพยาบาล (Infrastructure) และการเงิน (Finance) ตามลำดับ ส่วนโรงพยาบาลรัฐพบปัญหาสูงสุดในด้าน โครงสร้างโรงพยาบาล เจ้าหน้าที่ การดำเนินงาน การบริการ และการเงิน เมื่อเปรียบเทียบทั้งโรงพยาบาลสองประเภท จะสามารถเห็นได้ชัดเจนว่า ในด้านเจ้าหน้าที่ เป็นปัญหาที่พบมากในทั้งโรงพยาบาลรัฐและเอกชน แต่ในด้านของโครงสร้างโรงพยาบาล เห็นได้อย่างชัดเจนว่าเป็นปัญหาหลักของโรงพยาบาลรัฐที่ควรได้รับการแก้ไขและปรับปรุง ซึ่งปัญหานี้พบได้ค่อนข้างน้อยในโรงพยาบาลเอกชน ส่วนในด้านการบริการ และการดำเนินการและด้านการเงิน ไม่ค่อยมีความแตกต่างกัน แต่จะเห็นว่าด้านการเงินเป็นปัญหาที่พบน้อยสุดทั้งโรงพยาบาลรัฐเอกชน ซึ่งโดยทั่วไปแล้ว โรงพยาบาลเอกชนมีค่าบริการที่สูงกว่า อาจจะมีแนวโน้มในการพบปัญหานี้ได้มากกว่าโรงพยาบาลรัฐบาล แต่ปรากฏว่าปัญหาการเงินมีการถูกพูดถึงในเชิงลบน้อยที่สุดไม่แตกต่างกัน และในภาพรวมของการวิเคราะห์จะสามารถสรุปได้ว่า โรงพยาบาลรัฐถูกพูดถึงในเชิงลบมากกว่าโรงพยาบาลเอกชนเป็นส่วนใหญ่ นอกจากนั้นยังทำให้ทราบอีกว่า ผู้ใช้บริการให้ความสำคัญในเรื่องค่าใช้จ่ายน้อยกว่าด้านอื่นๆทั้งของโรงพยาบาลรัฐและเอกชนดังที่รายงานการวิเคราะห์บนแดชบอร์ด

บทที่ 5

สรุปผลการวิจัย อภิปรายผล และข้อเสนอแนะ

ในงานวิจัยนี้มีจุดมุ่งหมายเพื่อนำเสนอการรายงานผลการวิเคราะห์ความรู้สึกรวบรวมข้อมูลจากประสบการณ์การใช้บริการของผู้ป่วย ด้วยวิธีการใช้เทคนิคการเรียนรู้ของเครื่อง ด้วยการทำงานของโปรแกรมประมวลผลภาษาธรรมชาติ ตั้งแต่การพัฒนาโมเดลในเก็บข้อมูลอย่างอัตโนมัติของข้อมูลรีวิวในเว็บไซต์ honestdocs ผู้วิจัยได้ผลการเปรียบเทียบประสิทธิภาพของแบบจำลองด้วยอัลกอริทึมที่สนใจทั้ง 3 อัลกอริทึมและนำข้อมูลที่ได้มาสร้างการรายงานผลและวิเคราะห์หาข้อมูลเชิงลึก เพื่อนำสรุปผลโดยแบ่งหัวข้อการสรุปผลได้ดังนี้

- 1.สรุปผลการวิจัย
- 2.การอภิปรายผล
- 3.ข้อจำกัดและอุปสรรคของงานวิจัย
- 4.ข้อเสนอแนะ

5.1 สรุปผลการวิจัย

ในงานวิจัยนี้ได้ศึกษาการทำวิเคราะห์ความรู้สึกด้วยการใช้เทคนิคการเรียนรู้ของเครื่อง เพื่อสร้างแบบจำลองในการแยกประเภทความรู้สึกของข้อมูลรีวิวจากประสบการณ์การใช้บริการของผู้ป่วยในโรงพยาบาลในประเทศไทยที่ได้สกัดและเก็บรวบรวมมาจากเว็บไซต์ โดยข้อมูลที่น่าสนใจในงานวิจัยเป็นข้อมูลภายนอก และงานวิจัยนี้สามารถนำไปใช้ประโยชน์จากการใช้เทคนิคการสร้างแบบจำลองเป็นเบื้องหลังในการวิเคราะห์ข้อมูลเพื่อนำไปสร้างรายงานผลการวิเคราะห์ข้อมูลที่ได้จากการเก็บข้อมูลอย่างอัตโนมัติและการวิเคราะห์ที่รวดเร็วและนำไปสู่การรับรู้ปัญหาที่เกิดขึ้นของผู้บริหารหรือผู้กำหนดนโยบายในระบบการจัดการของโรงพยาบาลหรือแม้แต่ธุรกิจอื่นที่ได้มีการใช้ข้อมูลของผู้ใช้บริการมาวิเคราะห์เพื่อใช้ในการปรับปรุงธุรกิจให้ดีขึ้น

ปัจจัยและข้อจำกัดที่ส่งผลกระทบต่อการศึกษาของงานวิจัยนี้ คือ ข้อมูลรีวิวที่นำมาใช้ เนื่องจากเราดึงข้อมูลรีวิวจากเว็บไซต์ออกมาเป็นข้อความในระบบ Comment และการเตรียมข้อมูลในการฝึกฝนและทดสอบ จะทำให้ข้อมูลอยู่ในรูปแบบระดับ Sentence ดังนั้นเมื่อข้อความในระดับ Comment ถูกเลเบลด้วยคลาสใด ข้อความในระดับ Sentence จะถูกเลเบลไปด้วย เนื่องจากจุดประสงค์ของงานวิจัยต้องการจะทราบว่าผู้ใช้บริการพูดถึงเรื่องใดบ้างในข้อความรีวิว เพื่อนำมาแยกข้อความรีวิวเป็นหัวข้อตัวชี้วัดต่างๆ ซึ่งข้อมูลที่ได้ อาจจะมีข้อมูลที่ผิดพลาดจากการเลเบลไม่ตรงความหมาย จึงอาจจะทำให้ข้อมูลยังมีคุณภาพไม่สมบูรณ์ ความถูกต้องของข้อมูลซึ่ง

จำเป็นต้องให้ความสำคัญกับความถูกต้องของข้อมูลในขั้นตอนการเตรียมข้อมูลเพื่อให้ข้อมูลฝึกฝนและข้อมูลทดสอบที่มีคุณภาพมากที่สุด ซึ่งในขั้นตอนนี้ยังเป็นข้อจำกัดของงานวิจัยนี้

ในอนาคตสามารถประยุกต์ใช้งานวิจัยนี้ในการสร้างระบบอัตโนมัติสำหรับการวิเคราะห์ข้อมูลประสิทธิภาพการให้บริการได้โดยสร้างเป็นระบบ Support System เพื่อการเก็บรวบรวมข้อมูลและวิเคราะห์ เพื่อสนับสนุนการตัดสินใจแก้ไขปัญหาภายในโรงพยาบาลได้

5.2 การอภิปรายผล

เนื่องจากในงานวิจัยนี้ได้ผลการประเมินการเปรียบเทียบประสิทธิภาพแบบจำลองด้วยอัลกอริทึม Random Forest สูงสุด เมื่อเปรียบเทียบกับอัลกอริทึมอื่น จากการใช้โมเดล Random forest วิเคราะห์แยกประเภทความรู้สึกยังมีค่า Recall ไม่ดีเท่าที่ควร จึงได้ทำการตรวจสอบข้อมูลจากการดูค่า prediction probability ที่มีค่า probability ในการทำนายที่สูง ที่แสดงถึงความคิดเห็นเชิงบวก แต่ข้อมูลจริงระบุประเภทเป็น ความคิดเห็นเชิงลบ และจากการตรวจสอบความหมายของในของประโยคในข้อความรีวิว พบว่าประโยคย่อยส่วนใหญ่พูดไปในทางเชิงความหมายที่เป็นบวกจริง ยกตัวอย่างเช่น "โดยภาพรวมแล้วถือว่าดี ควรจัดแยกแผนกต้อนรับออกไปให้ชัดเจน" หรือ "พยาบาลบริการดี น่ารัก" จึงให้เหตุผลได้ว่า ในขั้นตอนในการดำเนินงานวิจัย ได้ดึงข้อมูลจากเว็บไซต์ออกมาเป็นข้อความรีวิวยาวๆ ที่ประกอบไปด้วยหลายประโยค จากนั้นนำมาแบ่งเป็นประโยคย่อย ซึ่งประโยคย่อยทุกประโยคจะถูกระบุประเภทความรู้สึกเดียวกันทั้งหมดของทั้งข้อความรีวิว เมื่อแบ่งแต่ละประโยคออกมา แต่ละประโยคจะสามารถมีทั้งประโยคที่บ่งบอกในความหมายเชิงบวกและเชิงลบ จึงทำให้เกิดการระบุประเภทข้อความรีวิวผิดพลาดได้ในกรณีที่ประโยคในข้อความรีวิวมีความกำกวม หรือมีบางประโยคที่ไม่ได้สื่อความหมายไปในเชิงเดียวกันทุกประโยคของข้อความรีวิว จึงมีทำให้ประโยคไม่มีความสอดคล้องกับประเภทความรู้สึกที่ถูกระบุไว้ ผู้วิจัยจึงทำข้อมูลการตรวจสอบคุณภาพของข้อมูลชุดทดสอบด้วยบุคคล โดยการสุ่มตรวจข้อมูลในแต่ละคลาส 50% พบว่าคุณภาพของข้อมูลติดฉลากถูกมีเพียง 43% ซึ่งสามารถนำมาเป็นเหตุผลหนึ่งที่ทำให้ประสิทธิภาพการทำนายโมเดลยังไม่ดีเท่าที่ควร โดยข้อมูลคุณภาพของชุดข้อมูลทดสอบแสดงดังตารางรายงานคุณภาพของข้อมูล

จากการใช้ข้อมูลรีวิวจากประสบการณ์การให้บริการของผู้ป่วยของโรงพยาบาลรัฐและเอกชนที่นำมาวิเคราะห์และศึกษาเพื่อหาข้อมูลเชิงลึก ผู้วิจัยได้ศึกษาข้อมูลเชิงลึกจากผลการวิเคราะห์ข้อมูลผ่านแดชบอร์ด เพื่อใช้ในการอภิปราย โดยการอภิปรายจะนำเสนอในमुखของผลจากการวิเคราะห์และในมุมมองที่เกิดขึ้นจริงกับคนปกติทั่วไป ดังนี้

ด้านเจ้าหน้าที่ : ในมุมมองของคนทั่วไปจะมองว่าหมอหรือเจ้าหน้าที่ของโรงพยาบาลเอกชนไม่เก่งเท่าโรงพยาบาลรัฐ เมื่อพิจารณาจากผลการวิเคราะห์แล้วพบว่าไม่มีความแตกต่างกันมาก เนื่องจากในโรงพยาบาลเอกชนบางครั้งอาจมีหมอจากโรงพยาบาลรัฐเข้าไปทำงานโรงพยาบาลเอกชนด้วย

ด้านการเงิน: ในมุมมองของคนทั่วไป อาจจะมองว่าโรงพยาบาลเอกชนมีค่าใช้จ่ายที่สูง ผู้ใช้บริการส่วนใหญ่จะเลือกมาใช้บริการโรงพยาบาลรัฐที่มีค่าใช้จ่ายที่ถูกลง แต่จากการวิเคราะห์ข้อมูลพบว่า ผู้ใช้บริการไม่ได้มีปัญหาทางการเงินเกี่ยวกับค่าใช้จ่ายที่เกิดขึ้นกับโรงพยาบาล เพราะด้านการเงินถูกพูดถึงในเชิงลบน้อยทั้งโรงพยาบาลรัฐและเอกชนพอๆกัน ซึ่งทำให้ทราบว่าผู้ให้บริการบางคนให้ความสำคัญกับการได้รับบริการที่ได้รับความสะดวกสบาย และการบริการที่ดีมากกว่าการเรื่องของค่าใช้จ่าย

ด้านโครงสร้าง: ในมุมมองของคนทั่วไป ตามหลักแล้วโรงพยาบาลเอกชนน่าจะมีโครงสร้างของโรงพยาบาลที่ดีกว่าโรงพยาบาลรัฐ จากผลการวิเคราะห์ข้อมูลพบว่าผู้ให้บริการจะถูกพูดถึงด้านโครงสร้างของโรงพยาบาลรัฐในเชิงลบมากกว่าเอกชนค่อนข้างมาก แต่ในความเป็นจริงยังมีโรงพยาบาลรัฐบางโรงพยาบาลที่มีโครงสร้างที่ดี เช่น โรงพยาบาลปิยมหาราณ หรือโรงพยาบาลที่อยู่ระดับพรีเมียม

ด้านการดำเนินงาน: ในมุมมองของคนทั่วไป จะมองว่าโรงพยาบาลเอกชนมีการดำเนินงานที่ไวกว่า สะดวกกว่าโรงพยาบาลรัฐ ซึ่งผลจากการวิเคราะห์ข้อมูลพบว่า โรงพยาบาลรัฐถูกพูดถึงในด้านการดำเนินงานในเชิงลบมากกว่าโรงพยาบาลเอกชน ซึ่งในสถานการณ์จริงปัจจุบันโรงพยาบาลเอกชนมีการให้บริการที่ดีและผู้ใช้บริการมีความสะดวกสบายมากกว่า ไม่ควรนอน

ด้านการบริการ: ในมุมมองของคนทั่วไป จะมองว่าเรื่องของการบริการของโรงพยาบาลรัฐเอกชนน่าจะดีกว่าโรงพยาบาลรัฐ แต่จากการวิเคราะห์ข้อมูล จะมองว่าโรงพยาบาลเอกชนมีการดำเนินงานที่ไวกว่า สะดวกกว่าโรงพยาบาลรัฐ ซึ่งผลจากการวิเคราะห์ข้อมูลพบว่าโรงพยาบาลรัฐเอกชนไม่มีความแตกต่างกันมาก ซึ่งในความเป็นจริง ทั้งโรงพยาบาลรัฐและเอกชนเมื่อมีช่วงที่วุ่นวายส่วนใหญ่มักเกิดที่ห้องฉุกเฉิน เจ้าหน้าที่อาจจะให้การบริการที่ไม่ดีจากสถานการณ์ที่วุ่นวาย

ด้านอื่นๆ: ในด้านอื่นจากข้อมูลแสดงผลการวิเคราะห์ เป็นการพูดถึงในเรื่องทั่วไปในเชิงลบ เช่น การบ่นถึงคนรอบข้างที่มาใช้บริการด้วยกัน หรือการบ่นโดยที่ไม่ได้สื่อถึงปัญหาที่เกี่ยวกับโรงพยาบาลทั้ง 5 ด้านตัวชี้วัด

5.3 ข้อจำกัดและอุปสรรคของงานวิจัย

เนื่องจากข้อมูลรีวิวนำมาใช้ฝึกฝนและทดสอบในการสร้างแบบจำลองการเรียนรู้ของเครื่องเป็นข้อความภาษาอังกฤษที่ได้มาจากการใช้ไลบรารี Google Translate แปลภาษาไทยให้เป็นภาษาอังกฤษ เพื่อให้ง่ายต่อการประมวลภาษาธรรมชาติก่อนนำเข้าแบบจำลองแต่ยังมีข้อจำกัดในเรื่องความถูกต้องในการแปลภาษา เนื่องจากข้อความภาษาไทยบางคำมีความกำกวม ความไม่ชัดเจน หรือเป็นภาษาที่แสดงความคิดเห็นแบบประชดประชัน จึงไม่ได้ให้ความหมายที่แท้จริง ทำให้ยังเกิดความผิดพลาดในความชัดเจนของภาษาได้และความถูกต้องของข้อมูลที่นำมาใช้ในการฝึกฝนและทดสอบกับแบบจำลอง และนอกจากนั้นในการศึกษางานวิจัยนี้ในการนำข้อมูลมาใช้ภายนอกจากการสกัดจากเว็บไซต์ และใช้การติดเลเบลข้อมูลจากคะแนน จากคะแนนความพึงพอใจในระดับความคิดเห็นมาใช้ในการเลเบลข้อมูลระดับประโยคด้วยจึงทำให้ข้อความในบางประโยคที่อาจจะไม่ได้มีความหมายตรงกับเลเบลที่ถูกติดไว้ ซึ่งมีผลต่อคุณภาพของข้อมูล เนื่องจากในงานวิจัยต้องการได้ข้อมูลในระดับลึกและรายละเอียดว่าผู้ที่มีประสบการณ์การใช้บริการพูดถึงโรงพยาบาลในด้านใดบ้าง ผู้วิจัยจึงใช้วิธีการนำข้อมูลในระดับประโยคเหล่านี้มาวิเคราะห์ให้ลึกขึ้น โดยยังคงมีข้อจำกัดในเรื่องของคุณภาพ จึงมีการตรวจสอบและแสดงผลรายงานจากการสุ่มข้อมูลเพื่อดูคุณภาพของข้อมูลประกอบกับการใช้ข้อมูลสำหรับนำเข้าแบบจำลอง ดังนั้นผู้วิจัยจึงจำเป็นต้องศึกษาวิธีการในการจัดเตรียมข้อมูลให้มีความถูกต้องและชัดเจนมากขึ้น เพื่อประสิทธิภาพที่ดีของข้อมูลและแบบจำลอง

5.4 ข้อเสนอแนะ

ด้านการปรับปรุงปริมาณข้อมูล เนื่องจากในงานวิจัยนี้ใช้ข้อมูลประเภทข้อความภาษาไทยที่เกี่ยวกับการแสดงความคิดเห็นในการใช้บริการของโรงพยาบาล จากแหล่งข้อมูล www.honestdocs.co.th เพียงแหล่งข้อมูลเดียว และข้อมูลค่อนข้างมีปริมาณน้อย หากเปรียบเทียบในแต่ละโรงพยาบาล ซึ่งยังมีแหล่งข้อมูลอื่นที่มีโครงสร้างของข้อความรวิวนีที่ประกอบไปด้วย ชื่อโรงพยาบาล ข้อความรีวิว และคะแนนรีวิว หรือมีองค์ประกอบข้อมูลคล้ายเคียงกับงานวิจัยนี้ บนแหล่งข้อมูลอื่นๆอีกเช่น การรีวิวบน google maps ที่มีโครงสร้างเหมือนกัน สามารถนำมาใช้ร่วมกัน กับชุดข้อมูลที่มีอยู่ได้เพื่อเพิ่มจำนวนข้อมูลตัวอย่างสำหรับการใช้ฝึกฝนโมเดลการเรียนรู้ของเครื่อง

1. ด้านของการปรับปรุงคุณภาพของข้อมูล บนงานวิจัยนี้พบว่าปัญหาที่ทำให้เกิดค่าความถูกต้องที่ได้จากการเรียนรู้ของเครื่องด้วยอัลกอริทึมต่างๆ มีประสิทธิภาพได้ไม่ดีเท่าที่ควรเกิดจากการให้คะแนนรีวิวที่ไม่ได้สื่อถึงความหมายของข้อความรวิวย่างชัดเจน ในการปรับปรุงคุณภาพ

สามารถเพิ่มการคัดกรองโดยคัดเลือกเฉพาะข้อมูลที่มีการให้คะแนนหนักแน่นไปในทางเชิงลบหรือเชิงบวกอย่างชัดเจน เช่นการคัดเลือกเฉพาะสัดส่วนที่ได้รับคะแนน1 ซึ่งบอกว่าเป็นเชิงลบ หรือคะแนน 5 ที่บ่งบอกว่าเป็นเชิงบวก

2. ด้านการจำแนกหัวข้อตัวชี้วัด จำเป็นที่จะต้องวิเคราะห์เพิ่มเติม จากกราฟพายจะเห็นได้ว่าข้อมูลส่วนใหญ่ถูกจำแนกเป็นหมวดหมู่อื่นๆ ซึ่งในการปรับปรุง ควรเพิ่มตัวกรองสำหรับหมวดหมู่ที่มีอยู่ให้มากขึ้น เช่น การเลือกเฉพาะเจาะจงในบางคำที่สื่อไปถึงหัวข้อตัวชี้วัดนั้นอย่างชัดเจน และการค้นหาคำตกหล่นสำหรับการกรองด้วยการตรวจสอบบนเวิร์ดคลาวด์ว่า คำส่วนใหญ่ที่ตกหล่นควรเป็นหัวข้อตัวชี้วัดใหม่ หรือควรเพิ่มในหัวข้อตัวชี้วัดที่มีอยู่

3. ด้านการปรับปรุงคุณภาพของโมเดลการเรียนรู้ เนื่องจากงานวิจัยนี้มีการเปรียบเทียบเพียง อัลกอริทึม random forest, multinomial Naive Bayes, LSTM ซึ่งสามารถใช้อัลกอริทึมอื่นๆ ได้เช่นกัน เพื่อการหาอัลกอริทึมที่เหมาะสมในการทำงานกับข้อมูลชุดนี้

4. ด้านของการปรับปรุงปริมาณข้อมูล เนื่องจากในงานวิจัยนี้ใช้ข้อมูลประเภทข้อความภาษาไทยที่เกี่ยวกับการแสดงความคิดเห็นในการใช้บริการของโรงพยาบาล จากแหล่งข้อมูล www.honestdocs.co.th เพียงแหล่งข้อมูลเดียว และข้อมูลค่อนข้างมีปริมาณน้อย หากเปรียบเทียบในแต่ละโรงพยาบาล ซึ่งยังมีแหล่งข้อมูลอื่นที่มีโครงสร้างของข้อความรีวิวที่ประกอบไปด้วย ชื่อโรงพยาบาล ข้อความรีวิว และคะแนนรีวิว หรือมีองค์ประกอบข้อมูลคล้ายเคียงกับงานวิจัยนี้ บนแหล่งข้อมูลอื่นๆอีกเช่น การรีวิวบน google maps ข้อมูลจากแพลตฟอร์มในเฟซบุ๊กของโรงพยาบาล และแพลตฟอร์มอื่นๆ ที่มีการพูดถึงโรงพยาบาล ที่เพิ่มจำนวนข้อมูลตัวอย่างสำหรับการใช้ฝึกฝนโมเดลการเรียนรู้ของเครื่อง

บรรณานุกรม

- Haruechaiyasak, C., Kongthon, A., Palingoon, P., & Trakultaweekoon, K. (2013). *S-Sense: A Sentiment Analysis Framework for Social Media Sensing*. Paper presented at the SocialNLP@IJCNLP.
- Hassan, S.-U., Aljohani, N. R., Idrees, N., Sarwar, R., Nawaz, R., Martínez-Cámara, E., . . . Herrera, F. (2020). Predicting literature's early impact with sentiment analysis in Twitter. *Knowledge-Based Systems*, 192, 105383.
doi:<https://doi.org/10.1016/j.knosys.2019.105383>
- Jianqiang, Z., Xiaolin, G., & Xuejun, Z. (2018). Deep Convolution Neural Networks for Twitter Sentiment Analysis. *IEEE Access*, 6, 23253-23260.
doi:10.1109/ACCESS.2017.2776930
- Sajeevan, A., & L. K, S. (2019). An enhanced approach for movie review analysis using deep learning techniques. 2019 *International Conference on Communication and Electronics Systems (ICCES)*, 1788-1794. doi:10.1109/ICCES45898.2019.9002043
- Sewalk, K. C., Tuli, G., Hswen, Y., Brownstein, J. S., & Hawkins, J. B. (2018). Using Twitter to Examine Web-Based Patient Experience Sentiments in the United States: Longitudinal Study. *Journal of medical Internet research*, 20(10), e10043-e10043.
doi:10.2196/10043. (Accession No. 30314959)
- Shah, D., Isah, H., & Zulkernine, F. (2018). Predicting the Effects of News Sentiments on the Stock Market. 2018 *IEEE International Conference on Big Data (Big Data)*, 4705-4708. doi:10.1109/BigData.2018.8621884
- Wicaksono, A. T., & Mariyah, S. (2019). Social Network Analysis of Health Development in Indonesia. 2019 *3rd International Conference on Informatics and Computational Sciences (ICICoS)*, 1-6. doi:10.1109/ICICoS48119.2019.8982482





รายละเอียดโค้ดการดึงข้อมูลจากเว็บไซต์

```

1 from selenium import webdriver
2 from selenium.common.exceptions import NoSuchElementException
3 from webdriver_manager.chrome import ChromeDriverManager
4 import time
5 from random import randint
6 from time import sleep
7 from requests import get
8 from bs4 import BeautifulSoup
9 import pandas as pd
10 import time

```

```

1 #get html content
2 def getHTMLContent(url):
3     try:
4         response = get(url)
5         html = BeautifulSoup(response.text, 'html.parser')
6     except Exception as e:
7         html = None
8     return html

```

```

1 hospital_link = ['https://www.honestdocs.co/hospitals/world-medical-center-hospital',
2                  'https://www.honestdocs.co/hospitals/thanyaburi-hospital',
3                  'https://www.honestdocs.co/hospitals/samitivej-sriracha-hospital',
4                  'https://www.honestdocs.co/hospitals/pathumthani-hospital',
5                  'https://www.honestdocs.co/hospitals/navamin-1-hospital',
6                  'https://www.honestdocs.co/hospitals/bangkruai-hospital',
7                  'https://www.honestdocs.co/hospitals/nakhon-ping-hospital',
8                  'https://www.honestdocs.co/hospitals/ban-fang-hospital',
9                  'https://www.honestdocs.co/hospitals/bangkok-hospital-phitsanulok',
10                 'https://www.honestdocs.co/hospitals/khelang-nakhon-ram-hospital', 'https://www.honestdocs.co/hospitals/lai

```

Web scraping

```

1 i = 0
2 below_3 = []
3 for url in hospital_link:
4     try:
5         scrape_data = pd.read_csv("scrape_df.csv")
6     except:
7         scrape_data = pd.DataFrame({"name": [],
8                                     "number": [],
9                                     "link": [] })
10
11 hospital_name = url.split('/')[1]
12 driver = webdriver.Chrome('chromedriver', options=chrome_options)
13 # driver.get(url)
14 # sleep(3)
15 zero_comment = False
16
17 while not zero_comment:
18     try:
19         driver.get(url)
20         sleep(3)
21     except:
22         get_button = driver.find_element_by_class_name('ask__link.ask__link--load-more.open-reviews.js-load-more')
23     except NoSuchElementException:
24         print('0 comment for ', url)
25         below_3.append(url)
26         zero_comment = True
27         break
28     while True:
29
30         get_button.click()
31         sleep(5)
32         before = driver.find_element_by_class_name('reviews')
33         before_html = before.get_attribute('innerHTML')
34
35         get_button.click()
36         sleep(5)
37         after = driver.find_element_by_class_name('reviews')
38         after_html = after.get_attribute('innerHTML')
39         sleep(5)
40
41         if before_html == after_html:
42             success = True
43             print('The end')
44             sleep(1)
45             break

```

```

52 # Scrape data from html source
53 html_content = driver.page_source
54 driver.close()
55 html = BeautifulSoup(html_content, 'html.parser')
56
57 ratings = []
58 reviews = []
59 dates = []
60
61 comments = html.find_all('div', class_='comments__container')
62
63 for comment in comments:
64
65     rating = comment.find("span", class_="star-rating")["data-score"]
66     ratings.append(rating)
67
68     review = comment.find("div", class_="comments__content").p.text
69     reviews.append(review)
70     date = comment.find('p', class_='comments__date').text
71     dates.append(date[1:-1])
72
73
74 df = pd.DataFrame({"reviews": reviews,
75                   "ratings": ratings,
76                   "dates": dates })
77 df['dates'] = pd.to_datetime(df['dates'], format="%B %d, %Y %H:%M")
78
79 df.to_csv("data/" + hospital_name + ".csv", index=False, encoding='utf-8-sig')
80
81 data = pd.DataFrame([[hospital_name, len(reviews), url]], columns=['name', 'number', 'link'])
82 scrape_data = pd.concat([scrape_data, data])
83 scrape_data.to_csv("scrape_data1.csv", index=False, encoding='utf-8-sig')
84 sleep(randint(1,5))
85 i += 1
86 print('index', i)
87 print(hospital_name, len(reviews))
88 print(below_3)

```



ภาควิชาคณิตศาสตร์

รายละเอียดได้การเตรียมข้อมูลด้วยเทคนิคการประมวลภาษาธรรมชาติ

```

1  #NLP
2  import nltk
3  from nltk import sent_tokenize, word_tokenize, RegexpParser
4  from nltk.sentiment.util import mark_negation
5  from nltk.corpus import stopwords, wordnet
6  from nltk.stem.wordnet import WordNetLemmatizer
7  from nltk.tag import pos_tag
8  nltk.download('punkt')
9  nltk.download('stopwords')
10 nltk.download('averaged_perceptron_tagger')
11 nltk.download('wordnet')
12
13 def plural2singular(word):
14     if word == "doctors":
15         return "doctor"
16     elif word == "nurses":
17         return "nurse"
18     elif word == "clinics":
19         return "clinic"
20     elif word == "hospitals":
21         return "hospital"
22     elif word == "services":
23         return "service"
24     elif word == "staffs":
25         return "staff"
26     elif word == "treatments":
27         return "treatment"
28     elif word == "students":
29         return "student"
30     else:
31         return word
32
33 def get_wordnet_pos(tag):
34     if tag.startswith('J'):
35         return wordnet.ADJ
36     elif tag.startswith('V'):
37         return wordnet.VERB
38     elif tag.startswith('N'):
39         return wordnet.NOUN
40     elif tag.startswith('R'):
41         return wordnet.ADV
42     else:
43         return 'a'
44
45 def lemmatize_word(tokenized_words):
46     lemmatizer = WordNetLemmatizer()
47     lemmatized_words = []
48     for word, tag in pos_tag(tokenized_words):
49         pos = get_wordnet_pos(tag)
50         lemmatized_words.append(lemmatizer.lemmatize(word, pos))
51     lemmatized_words = ' '.join([str(word) for word in lemmatized_words])
52     return lemmatized_words
53

```

```

54 def extract_noun_phrases(parsed_trees):
55     nps = []
56     for parsed_tree in parsed_trees:
57         np = []
58         for subtree in parsed_tree.subtrees():
59             if subtree.label() == 'NP':
60                 t = subtree
61                 t = " ".join(word for word, tag in t.leaves())
62                 np.append(t)
63         nps.append(np)
64     return nps
65
66 def preprocess_review(reviews):
67     cleaned_reviews = []
68     for review in reviews:
69         review = re.sub("[^a-zA-Z0-9\s]", " ", review)
70         words = word_tokenize(review)
71         tokenized_words = [(word.lower()).strip() for word in words if word not in stopwords.words("english")]
72         lemmatized_word = lemmatize_word(tokenized_words)
73
74         cleaned_reviews.append(lemmatized_word)
75
76 def create_parsed_trees(df):
77     # Define your custom grammar (modified to be a valid regex)
78     grammar = """NP: {<NN|NNS>+}
79             {<NN|NNS>+<CC>{<NN|NNS>+}}"""
80
81     # Create an instance of your custom parser
82     chunker = RegexpParser(grammar)
83     parse_trees = []
84     for sentence in df['reviews_en']:
85         parse_trees.append(chunker.parse(pos_tag(word_tokenize(sentence))))
86     return parse_trees

```

NLP Processing & Predicting

```
1 df = df[df.sentiment != "neutral"]
2 df['sentiment'] = df['sentiment'].replace('positive',0)
3 df['sentiment'] = df['sentiment'].replace('negative',1)

1 df['reviews_en'] = df['reviews_en'].apply(str)
2 df = extract_sentences(df)
3 df["cleaned_reviews"] = preprocess_review(df["reviews_en"])
4 df["sentiments_pred"] = predict_sentiment(multinomial_nb_classifier ,df)
5 df["NP"] = extract_noun_phrases(create_parsed_trees(df))
```





ภาคผนวก ค

รายละเอียดโค้ดการจำแนกประเภทความรู้สึกด้วยแบบจำลองการเรียนรู้ของเครื่อง

```

1 #Utils
2 import numpy as np
3 import pandas as pd
4 import matplotlib.pyplot as plt
5 import re
6 from tqdm import tqdm
7
8 #Machine Learning
9 from sklearn.model_selection import GridSearchCV, StratifiedKFold, train_test_split
10 from sklearn.feature_extraction.text import TfidfVectorizer, CountVectorizer
11 from sklearn.pipeline import Pipeline
12 from sklearn.naive_bayes import MultinomialNB
13 from sklearn.naive_bayes import MultinomialNB
14 from sklearn.ensemble import RandomForestClassifier
15 from sklearn.svm import SVC
16
17 # matrix classification
18 from sklearn.metrics import precision_recall_fscore_support
19 from sklearn.metrics import accuracy_score, ClassificationReport
20 # metric
21 from sklearn.metrics import classification_report
22
23
24 def train_random_forest(X_train, X_test, y_train, y_test):
25
26     tfidf_vectorizer = TfidfVectorizer()
27     rf_clf = RandomForestClassifier(random_state=2020, class_weight="balanced")
28     rf_pipeline = Pipeline([
29         ('tfidf', tfidf_vectorizer),
30         ('classifier', rf_clf)
31     ])
32     rf_params = {
33         "classifier__n_estimators": range(50, 200, 50),
34         "classifier__max_depth": range(2, 10, 2),
35         "classifier__min_samples_leaf": range(1, 4)
36     }
37     grid = GridSearchCV(rf_pipeline, param_grid=rf_params, cv=10, scoring='roc_auc')
38     grid.fit(X_train, y_train)
39
40     print("Best Score : %.2f" % grid.best_score_)
41     print("Best Params : ")
42     print(grid.best_params_)
43
44     y_pred = grid.predict(X_test)
45     print("Accuracy : %.2f" % accuracy_score(y_test, y_pred))
46     print(classification_report(y_test, y_pred))
47     return grid.best_estimator_
48
49
50 def predict_sentiment(estimator, df):
51     sentiments = []
52     for sentence in df["reviews_en"]:
53         sentiment = estimator.predict([sentence])
54         if sentiment == 1:
55             sentiment = "positive"
56         else:
57             sentiment = "negative"
58         sentiments.append( sentiment )
59     return sentiments

```

```

1 df = pd.read_csv("/content/gdrive/My Drive/Colab Notebooks/Thesis/train_test_bkk_translated.csv", index_col=False, header

```

```

1 train_test_data = df.iloc[: ]
2 train_test_data = train_test_data.reset_index()
3 train_test_data = extract_sentences(train_test_data)
4 train_test_data["cleaned_reviews"] = preprocess_review(train_test_data["reviews_en"])

```



รายละเอียดโค้ดการจัดหมวดหมู่ข้อความรีวิวตามหัวข้อตัวชี้วัด

```
1 df = df.loc[:,['reviews_en', 'cleaned_reviews', 'dates', 'sentiments_pred', 'NP', 'Hospital']]
2 df['staff_NP'] = ""
3 df['service_NP'] = ""
4 df['finance_NP'] = ""
5 df['infrastructure_NP'] = ""
6 df['process_NP'] = ""
7 df['other'] = ""
8 df['catagories'] = ""
9
10 Dimstaff = ["doctor", "nurse", "student", "guard", "staff", "Doctor", "Nurse"]
11 Diminfrastructure = ["hospital", "clinic", "equipment", "place", "building", "room", "department", "floor"]
12 Dimservice = ["care", "treatment", "advice", "surgery", "disease", "medical care", "service", "responsibility", "manner",
13 Dimfinance = ["price", "claim", "payment", "cost", "finance"]
14 Dimprocess = ["queue", "appointment", "process", "management", "schedule", 'Appointment']
15
16 for index, row in df.iterrows():
17     staff = []
18     service = []
19     finance = []
20     infrastructure = []
21     process = []
22     other = []
23     catagory = []
24
25     for i in range(len(row['NP'])):
26         # print(row['NP'][i])
27         if [word for word in Dimstaff if word in row['NP'][i]]:
28             staff.append(row['NP'][i])
29         elif [word for word in Diminfrastructure if word in row['NP'][i]]:
30             infrastructure.append(row['NP'][i])
31         elif [word for word in Dimservice if word in row['NP'][i]]:
32             service.append(row['NP'][i])
33         elif [word for word in Dimfinance if word in row['NP'][i]]:
34             finance.append(row['NP'][i])
35         elif [word for word in Dimprocess if word in row['NP'][i]]:
36             process.append(row['NP'][i])
37         else:
38             other.append(row['NP'][i])
39
40     if staff:
41         catagory.append("staff")
42     if infrastructure:
43         catagory.append("infrastructure")
44     if service:
45         catagory.append("service")
46     if finance:
47         catagory.append("finance")
48     if process:
49         catagory.append("process")
50     if not staff and not infrastructure and not service and not finance and not process:
51         catagory.append("other")
52
53     # Change List to string
54     staff = ' '.join(e for e in staff)
55     service = ' '.join(e for e in service)
56     finance = ' '.join(e for e in finance)
57     infrastructure = ' '.join(e for e in infrastructure)
58     process = ' '.join(e for e in process)
59     other = ' '.join(e for e in other)
60     catagory = ' '.join(e for e in catagory)
61
62     df['NP'][index] = ' '.join(e for e in row['NP'])
63     df['staff_NP'][index] = staff
64     df['infrastructure_NP'][index] = infrastructure
65     df['service_NP'][index] = service
66     df['finance_NP'][index] = finance
67     df['process_NP'][index] = process
68     df['other'][index] = other
69     df['catagories'][index] = catagory
```

ประวัติผู้เขียน

ชื่อ-สกุล	พนิดา กาศกลางดอน
วัน เดือน ปี เกิด	24 มกราคม 2538
สถานที่เกิด	จังหวัดอุทัยธานี
วุฒิการศึกษา	ปริญญาตรี หลักสูตรวิทยาศาสตร์บัณฑิต สาขาวิชาเคมี มหาวิทยาลัยศรีนครินทรวิโรฒ
ที่อยู่ปัจจุบัน	87 หมู่ 7 ตำบลห้วยแห้ง อำเภอบ้านไร่ จังหวัดอุทัยธานี 61140

