



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

การพัฒนาระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ของนักศึกษานบนเพชบุ๊ก

โดยใช้ตัวจำแนกข้อมูล นาอ์ฟ เบย์ สำหรับภาษาไทย

Development of a Real-time Sentiment Analysis System of Students

on Facebook Using Naive Bayes Classifier in Thai Language

สมักร ชัยสงวน (Samark Chaisanguan)¹ วณิชยา ร่มสายหยุด (Walisa Romsaiyud)²

บทคัดย่อ

การวิจัยครั้งนี้มีวัตถุประสงค์ 1) เพื่อพัฒนาระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ของนักศึกษานบนเพชบุ๊กโดยใช้ตัวจำแนกข้อมูลนาอ์ฟ เบย์ สำหรับภาษาไทย และ 2) เพื่อประเมินมาตรการความถูกต้องของค่าความแม่นยำและค่าเรียกคืนของระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ของนักศึกษานบนเพชบุ๊กโดยใช้ตัวจำแนกข้อมูลนาอ์ฟ เบย์ สำหรับภาษาไทย ดำเนินการโดยรวบรวมข้อมูลจากเพชบุ๊กเพจของมหาวิทยาลัยสุโขทัยธรรมาธิราช ซึ่งเป็นข้อมูลการแสดงความรู้สึกของนักศึกษาที่ติดตามเพชบุ๊กของมหาวิทยาลัย เพื่อทำการวิเคราะห์ความรู้สึก โดยหลักทฤษฎีการประมวลผลภาษาธรรมชาติด้วยวิธีการตัดคำภาษาไทย และวิเคราะห์ความรู้สึกโดยใช้เทคนิคการจำแนกข้อมูลนาอ์ฟ เบย์ ซึ่งสามารถแสดงผลแบบเรียลไทม์เช่นระบบจัดการคำคลังคำศัพท์ สรุปผลการวิเคราะห์ความรู้สึก เครื่องมือทดสอบประโยค และระบบแจ้งเตือนข้อความผ่านทางไลน์ ผลการวิจัยนี้ได้ผลลัพธ์ค่าความแม่นยำร้อยละ 96.61 ค่าเรียกคืนร้อยละ 96.50 ค่าความถูกต้อง 97.60 และการวัดประสิทธิภาพโดยรวมร้อยละ 96.55 ตามลำดับ

คำสำคัญ การวิเคราะห์ความรู้สึก จำแนกข้อมูลนาอ์ฟ เบย์ ระบบเรียลไทม์ เพชบุ๊กเพจของมหาวิทยาลัยสุโขทัยธรรมาธิราช

¹ นักศึกษาหลักสูตรปริญญาโท สาขาวิชาวิทยาศาสตร์และเทคโนโลยี แขนงเทคโนโลยีสารสนเทศและการสื่อสาร
มหาวิทยาลัยสุโขทัยธรรมาธิราช Samark.cha@stou.ac.th

² ผู้ช่วยศาสตราจารย์ ดร. ประจักษ์สาขาวิชาวิทยาศาสตร์และเทคโนโลยี แขนงเทคโนโลยีสารสนเทศและการสื่อสาร
มหาวิทยาลัยสุโขทัยธรรมาธิราช Walisa.rom@stou.ac.th



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

Abstract

The purposes of this research were as follows: 1) to develop a real-time sentiment analysis system of students on Facebook using Naive Bayes classifier in Thai language, and 2) to evaluate the accuracy of precision and recall measures of real-time sentiment analysis system of students on Facebook using Naive Bayes classifier in Thai language. The research collected data from the Facebook page of Sukhothai Thammathirat Open University; the data comprised the sentiment expression of students who followed the university's Facebook fan page. This method was employed as the means to perform an analysis of students' sentiment using Natural Language Processing with a Thai word segmentation method and Naive Bayes Classifier, All results were able to demonstrate in real-time results such as vocabulary management system, the results of sentiment analysis, the tool for testing sentences, and Line notification system. The measurement results were 96.61% for precision, 96.50% for recall, 97.60% for accuracy and 96.55% for F-measure respectively.



Keywords: Sentiment analysis, Naive bayes classifier, Real-time system, STOU Facebook Fan page



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

บทนำ

เฟซบุ๊กเป็นโซเชียลมีเดียยอดนิยมอันดับ 1 ของคนไทย โดยมีจำนวนผู้ใช้งานอยู่ที่ 51 ล้านคน คิดเป็นร้อยละ 90 ซึ่งจากสถิติประเทศไทยยังเป็นประเทศที่ใช้เวลาต่อวันอยู่กับอินเทอร์เน็ตมากที่สุดในโลก โดยเฉลี่ยอยู่ที่ 9 ชั่วโมง 38 นาทีต่อวัน และใช้เวลาในการเล่นโซเชียลมีเดียโดยเฉลี่ยอยู่ที่ 3 ชั่วโมง 10 นาทีต่อวัน (wp, 2018) ด้วยเฟซบุ๊กถือว่าเป็นเครื่องมือที่นิยมใช้ในการติดต่อสื่อสาร และเป็นแหล่งแสดงความรู้สึกของผู้คนจำนวนมากการนำข้อมูลจากเพจเฟซบุ๊กมาใช้ในการวิเคราะห์ความรู้สึก (Sentiment Analysis) ของผู้ติดตามเพจในเฟซบุ๊กนั้นถือว่าสำคัญ โดยเฉพาะในด้านของการประสบความสำเร็จทางธุรกิจเพราะ “ลูกค้า” คือปัจจัยสำคัญที่สุดที่จะทำให้ธุรกิจประสบความสำเร็จ การสร้างความสัมพันธ์ที่ดีเพื่อรักษารฐานข้อมูลลูกค้าคือการ “รับความรู้สึกของลูกค้า” เป็นกระบวนการทางธุรกิจในการเข้าใจลูกค้า เพื่อใช้ประโยชน์ในการพัฒนาสินค้าและบริการ โดยส่งเสริมให้เกิดการจูงใจที่ดีต่อสินค้าและผลิตภัณฑ์ นอกจากนี้หน่วยงานภาครัฐสามารถนำความรู้ไปดำเนินการประยุกต์ในด้าน “ความรู้สึกของประชาชน” กับหลายๆ หน่วยงานภาครัฐ ซึ่งเป็นผลให้ข้อมูลมีปริมาณมากและเพิ่มขึ้นอยู่ตลอดเวลาอย่างรวดเร็ว เป็นเรื่องยากที่จะรวบรวมข้อมูลทั้งหมดมาอ่านหรือวิเคราะห์โดยการใช้อย่างมีประสิทธิภาพจึงจำเป็นต้องนำระบบคอมพิวเตอร์เข้ามาช่วยในการแก้ปัญหาเพื่อเพิ่มความสามารถในการแข่งขัน

ในปัจจุบันการศึกษาระดับอุดมศึกษามีการแข่งขันกันมากขึ้นผู้เรียนมีทางเลือกในการตัดสินใจเลือกสถานศึกษามากขึ้นมหาวิทยาลัยจึงจำเป็นต้องอย่างยิ่งที่จะต้องมีการพัฒนาระบบในการหาวิธีการที่จะให้กลุ่มเป้าหมายตัดสินใจเลือกเข้าศึกษาต่อในมหาวิทยาลัยโดยการจัดการกิจกรรมต่าง ๆ และรูปแบบการประชาสัมพันธ์ที่หลากหลายช่องทาง มหาวิทยาลัยจำเป็นต้องสร้างความพึงพอใจแก่ผู้เรียนหรือผู้มารับบริการด้วย มีความจำเป็นต้องอย่างยิ่งที่จะต้องรับรู้และเข้าใจความรู้สึกของนักศึกษาที่มีต่อมหาวิทยาลัย เพื่อมหาวิทยาลัยจะได้นำไปใช้เป็นแนวทางในการปรับปรุงและพัฒนาระบบการจัดการเรียนการสอนและ พัฒนาระบบการบริการของมหาวิทยาลัยต่อไป มหาวิทยาลัยสุโขทัยธรรมาธิราชเป็นหนึ่งในมหาวิทยาลัยในประเทศไทย ที่มีการใช้เฟซบุ๊กเป็นช่องทางในการติดต่อสื่อสารกับนักศึกษา โดยใช้ชื่อว่า STOUUniversity (<https://fb.com/STOUUniversity/>) มีนักศึกษาที่ติดตามเฟซบุ๊กเพจของมหาวิทยาลัยจำนวน 241,935 คน ซึ่งมีหลากหลายความคิดเห็นที่ส่งเข้ามาในเฟซบุ๊กเพจ ตัวอย่างความคิดเห็นเชิงบวกเช่น “ภูมิใจที่ได้เรียนมสธ” ความคิดเห็นเชิงลบเช่น “ข้อสอบยากมาก” และความคิดเห็นที่เป็นกลาง “แชร์ได้ไหมครับท่านอาจารย์”

ภาษาไทยเป็นภาษาที่มีลักษณะเฉพาะทางโครงสร้างที่ซับซ้อน ตั้งแต่โครงสร้างตัวอักษร ที่มีพยัญชนะ³ 44 รูป สระ 21 รูป วรรณยุกต์ 4 รูป ตัวเลข 10 รูป และเครื่องหมายวรรคตอนพร้อมทั้งสัญลักษณ์พิเศษต่าง ๆ เช่น ไม้ทัณฑฆาต จุลภาค และมหัพภาค เป็นต้น ภาษาไทยจึงเป็นภาษาที่ยากเมื่อเทียบกับภาษาอื่น เพราะเป็นภาษาที่ไม่มีการเขียนแบ่งพยางค์ คำ วลี หรือประโยค ไม่มีหลักเกณฑ์ตายตัว การสะกดคำมี กฎ ไวยากรณ์ ที่มีรูปแบบซับซ้อน มีคำยืม คำทับศัพท์ คำเฉพาะ คำเกิดใหม่จำนวนมาก และรวมถึงลักษณะของคำที่มีความกำกวมสูง จึงถือได้ว่า ประโยคในภาษาไทยมีความซับซ้อน ทำให้บางครั้งการตีความหมายของประโยคไม่เหมือนกัน แต่มีความหมายเหมือนกัน เช่น คำกำกวม ที่มีความหมายไม่แน่ชัด และภาษาไทยเป็นภาษาที่เรียงคำ การเรียงคำในภาษาไทยเป็นเรื่องที่สำคัญมาก เพราะการเปลี่ยนตำแหน่งของคำ ทำให้ความหมายของคำเปลี่ยนไปด้วย ซึ่งภาษาไทยไม่มีสัญลักษณ์ที่สามารถบ่งบอกถึงขอบเขตของคำ ที่เรียงต่อเนื่องกันทั้งประโยคเหมือนกับ

³ พยัญชนะหมายถึงเสียงพูดที่เปล่งออกมาโดยใช้อวัยวะส่วนต่าง ๆ ในปากและคอ เช่น เสียง ป โดยทั่วไปจะออกเสียงพยัญชนะร่วมกับเสียงสระ



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

ภาษาอังกฤษที่ใช้ช่องว่าง (Space) คั่นระหว่างขอบเขตของคำ เนื่องจากภาษาไทยเป็นภาษาที่ซับซ้อนจึงทำให้โปรแกรมคอมพิวเตอร์รู้จักคำในภาษาไทยนั้นมีความยุ่งยากกว่าในภาษาอังกฤษ โดยเฉพาะประโยคส่วนใหญ่บนอินเทอร์เน็ต ที่นิยมใช้ภาษาที่ไม่เป็นทางการหรือภาษาพูด รวมทั้งมีโครงสร้างประโยคที่ไม่ถูกต้องตามหลักไวยากรณ์ภาษาไทย ยิ่งทำให้ยากต่อการนำมาวิเคราะห์

งานวิจัยส่วนใหญ่มุ่งเน้นในการพัฒนาระบบการวิเคราะห์ความรู้สึกที่เป็นข้อความภาษาไทย ที่วิเคราะห์ความรู้สึกจากเฟซบุ๊กเพจที่เป็นข้อความภาษาไทยได้อย่างมีประสิทธิภาพและสามารถวิเคราะห์ได้แบบเรียลไทม์ งานวิจัยนี้จึงนำเสนอการพัฒนาระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time) ของนักศึกษานานาชาติ โดยใช้ตัวจำแนกข้อมูล นาอูฟ เบย์ สำหรับภาษาไทย ซึ่งเป็นวิธีจำแนกประเภทข้อมูลที่มีประสิทธิภาพวิธีหนึ่ง ซึ่งมีการนำไปประยุกต์ใช้งานในด้านการจำแนกประเภทข้อความ ที่ใช้งานได้ดีไม่ต่างจากการจำแนกประเภทวิธีการอื่น เนื่องจากเป็นวิธีการจำแนกข้อมูลได้อย่างมีประสิทธิภาพและอัลกอริทึมในการทำงานที่ไม่ซับซ้อน โดยผู้วิจัยจะรวบรวมข้อมูลโดยใช้ Facebook Graph API ร่วมกับ ภาษา PHP และ ภาษา Java Script (Socket.IO) มาเก็บไว้ที่ฐานข้อมูล เพื่อนำมาวิเคราะห์ความเห็นออกมาเป็นคะแนนบวก (+1) ลบ (-1) หรือเป็นกลาง (0) เพื่อวิเคราะห์ความรู้สึก (Sentiment Analysis) ของผู้ติดตามเพจที่เป็นเชิงบวก (Positive) และเชิงลบ (Negative)

วัตถุประสงค์

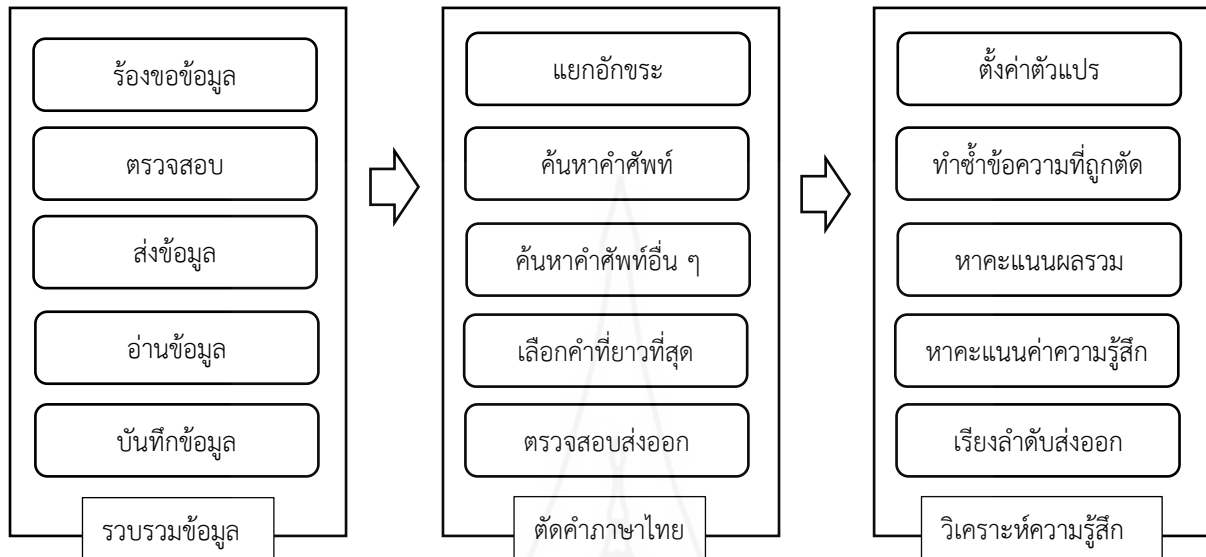
1. เพื่อพัฒนาระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ของนักศึกษานานาชาติเฟซบุ๊กโดยใช้ตัวจำแนกข้อมูลนาอูฟ เบย์ สำหรับภาษาไทย
2. เพื่อประเมินมาตรการความถูกต้องของค่าความแม่นยำและค่าเรียกคืนของระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ของนักศึกษานานาชาติเฟซบุ๊กโดยใช้ตัวจำแนกข้อมูลนาอูฟ เบย์ สำหรับภาษาไทย

ระเบียบวิธีวิจัย

วิธีการดำเนินงานของงานวิจัยนี้มุ่งเน้นเพื่อวิเคราะห์ออกแบบและพัฒนาระบบการวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time) สำหรับข้อความภาษาไทยเกี่ยวกับความคิดเห็นของนักศึกษาจากเฟซบุ๊กเพจ โดยมีกระบวนการทำงานของระบบวิเคราะห์ความรู้สึกเรียลไทม์ (Real-Time) ของนักศึกษานานาชาติเฟซบุ๊กโดยใช้ตัวจำแนกข้อมูลนาอูฟ เบย์ สำหรับภาษาไทย ดังภาพที่ 1



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8
The 8th STOU National Research Conference



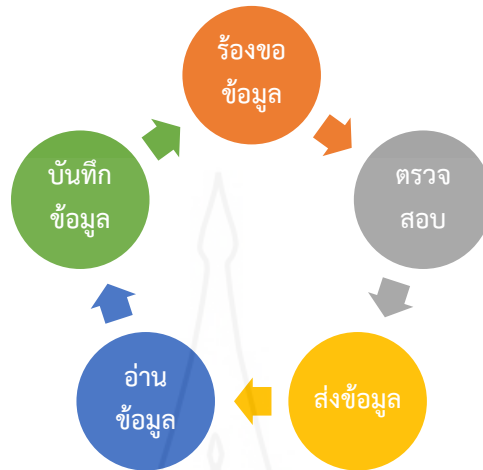
ภาพที่ 1 กระบวนการทำงานของระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time)

1. ข้อมูลและกลุ่มตัวอย่างที่ใช้ในงานวิจัย ระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time) ของนักศึกษานบนเฟซบุ๊ก โดยใช้ตัวจำแนกข้อมูลนาอ์ฟ เบย์ สำหรับภาษาไทยใช้ข้อมูลจากเฟซบุ๊กเพจมหาวิทยาลัยสุโขทัยธรรมาธิราช STOUUniversity (<https://fb.com/STOUUniversity/>) ซึ่งเป็นข้อมูลการแสดงความคิดเห็นของนักศึกษาที่ติดตามเฟซบุ๊กเพจของมหาวิทยาลัย จำนวน 241,935 คน และมีจำนวนข้อความที่ใช้ในการทดลองจำนวน 12,548 ข้อความ

2. กระบวนการทำงานของระบบในการรวบรวมข้อมูล พนิดา ทรงรัมย์ ได้ศึกษาการจำแนกความคิดเห็นที่มีต่อรัฐประหารในประเทศไทยจากความคิดเห็นบนเฟซบุ๊ก โดยวิธีการรวบรวมข้อมูลด้วย Facebook Graph API (พนิดา ทรงรัมย์, 2559) ซึ่งเป็นวิธีการเรียกใช้ข้อมูลจากเฟซบุ๊กบนอินเทอร์เน็ต (Internet) และเป็นวิธีการเดียวในการเรียกใช้ข้อมูลด้วยระบบคอมพิวเตอร์ งานวิจัยชิ้นนี้จึงใช้วิธีดังกล่าวในการรวบรวมข้อมูล โดยเริ่มจากร้องขอข้อมูลไปยังเฟซบุ๊ก จากนั้นเฟซบุ๊กจะทำการตรวจสอบว่ามีสิทธิ์ใช้งานข้อมูลหรือไม่ หากผ่านการตรวจสอบ เฟซบุ๊กจะส่งข้อมูลกลับมาที่เซิร์ฟเวอร์ (Server) หรือเครื่องแม่ข่าย ด้วยรูปแบบ JSON (JavaScript Object Notation) ระบบจะทำการอ่านและแปลงข้อมูลเพื่อบันทึกลงในฐานข้อมูลดังภาพที่ 2

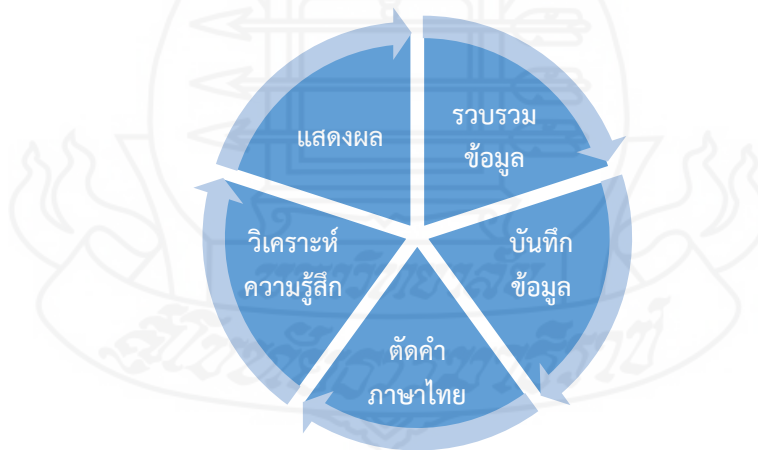


การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8
The 8th STOU National Research Conference



ภาพที่ 2 กระบวนการทำงานของระบบในการรวบรวมข้อมูล

3. กระบวนการทำงานโดยรวมของระบบ เมื่อได้รับข้อมูลจากเฟซบุ๊ก ระบบจะทำการบันทึกข้อมูลลงในฐานข้อมูล จากนั้นจะนำข้อมูลที่ได้มาเข้าสู่กระบวนการตัดคำภาษาไทย แล้วส่งข้อความที่ถูกตัดคำแล้วไปยังกระบวนการวิเคราะห์ความรู้สึก เพื่อระบุว่าข้อความดังกล่าวเป็นความรู้สึกประเภทใด หลังจากนั้นระบบจะส่งข้อมูลไปแสดงผลผ่านทางเว็บไซต์ (Web Site) ดังภาพที่ 3



ภาพที่ 3 กระบวนการทำงานโดยรวมของระบบ

4. กระบวนการตัดคำภาษาไทย ใช้การจับคู่ที่ยาวที่สุดและการจับคู่สูงสุด (Longest matching and maximal) โดยจะพิจารณาจากการวนลูปจากซ้ายไปขวาตามหลักภาษาไทยเพื่อเทียบกับพจนานุกรม Lexicon – Thai สามารถอธิบายด้วยอัลกอริทึม (Algorithm) การตัดคำภาษาไทยโดยมีขั้นตอนดังต่อไปนี้



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8 The 8th STOU National Research Conference

ขั้นตอนที่ 1 แยกอักขระ เมื่อนำข้อความเข้าสู่กระบวนการตัดคำสิ่งแรกที่ทำคือแยกอักขระทุกตัวออกจากกันแล้วจัดเก็บเข้าไปในตัวแปรชนิดอาร์เรย์ จากนั้นโปรแกรมจะนับว่ามีอักขระทั้งหมดกี่ตัวเพื่อหาความยาว แต่ตัวแปรชนิดอาร์เรย์เริ่มนับตำแหน่งจากศูนย์ จึงจำเป็นต้องลบด้วย 1 เพื่อระบุจำนวนตำแหน่งที่ถูกต้องทั้งหมดของอักขระ

ขั้นตอนที่ 2 ค้นหาคำศัพท์ เริ่มทำงานซ้ำจนกว่าค่าของตัวแปรระบุตำแหน่ง จะมีค่ามากกว่าตัวแปรความยาวของอักขระ ในแต่ละรอบจะนำข้อความไปเทียบกับพจนานุกรม หากพบข้อความจะเก็บไว้ในตัวแปรค่าที่พบในข้อความพร้อมกับระบุตำแหน่งที่พบ หากไม่พบจะนำอักขระต่อกับอักขระก่อนหน้าแล้วทำการเทียบกับพจนานุกรมอีกครั้ง

ขั้นตอนที่ 3 ค้นหาคำศัพท์อื่น ๆ กรณีที่พบคำศัพท์ในอักขระ โปรแกรมจะทำงานซ้ำ (For loop) โดยเริ่มจากตำแหน่งของอักขระปัจจุบันไปจนถึงตำแหน่งอักขระสุดท้าย โดยแต่ละรอบจะนำข้อความไปเทียบกับพจนานุกรมหากพบข้อความจะเก็บไว้ในตัวแปรค่าที่พบในข้อความ พร้อมกับระบุตำแหน่งที่พบ เพื่อแยกค่าที่พบออกจากกัน

ขั้นตอนที่ 4 เลือกคำที่ยาวที่สุด ถ้าโปรแกรมพบว่าตัวแปรค่าที่พบในข้อความไม่ใช่ค่าว่าง โปรแกรมจะเลือกคำศัพท์ที่ยาวที่สุดออกมา ด้วยหลักการอย่างง่ายคือ เลือกตำแหน่งของอาร์เรย์ (Array) ในลำดับสุดท้ายที่พบคำศัพท์ออกมา จะได้ข้อความที่ยาวที่สุดในแต่ละรอบแล้วเก็บข้อความที่ได้ไว้ในตัวแปรผลลัพธ์การตัดคำ พร้อมทั้งตั้งค่าตัวแปรค่าที่พบและตัวแปรข้อความที่พบชนิดอาร์เรย์ (Array) เป็นค่าเริ่มต้น หลังจากนั้นโปรแกรมจะทำซ้ำ (While loop) ข้อความในลำดับถัดไปโดยเริ่มจากตำแหน่งที่พบข้อความที่มากที่สุดจนกว่าค่าระบุตำแหน่ง (Pointer) จะมีค่ามากกว่าความยาวของอักขระ

ขั้นตอนที่ 5 ตรวจสอบและส่งออก ในกรณีที่ไม่มีพบคำศัพท์ใด ๆ โปรแกรมจะนำข้อความที่ไม่สามารถตัดได้ ไปเป็นผลลัพธ์ของการตัดคำ แต่หากพบคำศัพท์ ข้อความที่ถูกตัดจะอยู่ในรูปแบบของอาร์เรย์ ซึ่งเป็นค่าที่ถูกตัดเรียบร้อยแล้วเพื่อส่งออกไปทำงานในกระบวนการถัดไป สามารถเขียนเป็นแผนภาพรหัสเทียม (Pseudo code) ได้ดังภาพที่ 4

```
Input: Sentences
Output: Word Token
Read Sentences; Strings = Extract(Sentences)
While(Pointer=0; Pointer <= CountStringLength(Strings) -1)
    TempWord += String[Pointer] + TempWord
    If(TempWord === Dictionary's word)
        FoundWord[word] = TempWord
        For(pointer +1 > CountStringLength(Strings) -1; Pointer ++ )
            MoreStringTemp += MoreStringTemp + String[Pointer]
            If(MoreStringTemp === Dictionary's word)
                FoundWord[word] = MoreStringTemp
        ResultToken[] += Max(FoundWord)
    Pointer++
Return ResultToken
```

ภาพที่ 4 ขั้นตอนกระบวนการตัดคำออกจากข้อความภาษาไทย



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

จากภาพที่ 4 เป็นกระบวนการตัดคำภาษาไทยซึ่งสอดคล้องกับ 5 ขั้นตอนที่กำลังกล่าวไว้ข้างต้น โดยมีข้อมูลที่นำเข้าคือข้อความความคิดเห็น และผลลัพธ์ที่ได้คือ ข้อความที่ถูกตัดเรียบเรียงที่อยู่ในรูปแบบของข้อมูลประเภทแถวลำดับ (Array) เพื่อส่งไปขั้นตอนถัดไป

5. การจำแนกข้อมูลโดยใช้ตัวจำแนกนาอีย เบย์ (Naive Bayes Classifier) โดยการให้น้ำหนักคำในประโยคด้วยวิธีการเทียบค่าที่ถูกตัดกับฐานข้อมูลประเภทของคำ (Daniel Jurafsky and James H. Martin, 2018) ซึ่งมี 5 ประเภทด้วยกัน ในแต่ละประเภทของคำมีประเภทที่สนใจนำมาคิดค่าคะแนน 3 ประเภทคือคำที่สถานะเป็นกลาง, คำที่มีสถานะเชิงลบ และคำที่มีสถานะเชิงบวก ดังตารางที่ 1

ตารางที่ 1 ประเภทของคำในพจนานุกรม

ตัวอย่างคำ	ประเภท	รายละเอียด
นางสาว	Prefix	คำนำหน้าชื่อ
และแล้ว	Ignore	คำที่ไม่สนใจ
สังคม	Neutral	คำที่มีสถานะเป็นกลาง
ดีมาก ๆ	Positive	คำที่มีสถานะเป็นเชิงบวก
เลวทราม	Negative	คำที่มีสถานะเป็นเชิงลบ

เมื่อมีการตัดคำและกำหนดประเภทของคำภายในพจนานุกรมแล้วจะมีการจำแนกข้อมูลโดยการให้น้ำหนักของคำภายในประโยค เพื่อนำมาวิเคราะห์ความเห็นออกมาเป็นคะแนน เพื่อวิเคราะห์ความรู้สึก (นริศร พรหมบุตร, 2550) (นภดล สิทธิเลิศ, 2556) แบบ 3 กลุ่ม คือ ความรู้สึกเชิงบวก (Positive) ความรู้สึกเชิงลบ (Negative) และเป็นกลาง สามารถอธิบายด้วยอัลกอริทึม (Algorithm) การวิเคราะห์ความรู้สึก ด้วยทฤษฎีนาอีย เบย์ (Naive Bayes Classifier) โดยมีขั้นตอนแยกย่อยดังนี้

ขั้นตอนที่ 1 ตั้งค่าตัวแปร โดยแต่ละตัวแปรจะมีค่าดังนี้ ตัวแปร Classes มีค่าด้วยกัน 3 ค่าคือ positive, negative, neutral ตัวแปรค่าความน่าจะเป็นของแต่ละประเภทความรู้สึก (Prior) มีค่าเท่ากับ 0.3333 เท่ากัน เพราะความน่าจะเป็นของแต่ละเหตุการณ์มีค่าเท่ากับ 1/3 เท่ากับ 0.3333 และค่าคะแนนของแต่ละประเภทความรู้สึกมีค่าเท่ากับ 1

ขั้นตอนที่ 2 ทำซ้ำข้อความที่ถูกตัด นำคำที่แยกออกมาจากข้อความไปเทียบกับคำในตารางประเภทของคำ หากพบให้นำคำที่พบไปตรวจสอบว่าอยู่ในข้อความประเภทใด แล้วบวกค่าคะแนนของประเภทความรู้สึกนั้นด้วย 1 และเก็บไว้ที่ตัวแปรค่าคะแนนของแต่ละประเภทความรู้สึก (Score) ของคลาสนั้น ๆ ทำจนกว่าจะครบทุกคำในประโยค หลังจากนั้นนำค่าคะแนนของแต่ละคลาสไปคูณกับค่าความน่าจะเป็นของแต่ละประเภทความรู้สึก (Prior) และเก็บไว้ที่ตัวแปร (Score)

ขั้นตอนที่ 3 หาค่าคะแนนผลรวม ค่าความน่าจะเป็นของข้อความมีโอกาสเกิดขึ้นได้เพียงแค่ 1 แต่ในขั้นตอนที่ 2 นั้น หากพบข้อความจำนวนมาก จะทำให้ค่าคะแนนของแต่ละประเภทความรู้สึก (Score) มีค่ามากกว่า 1 ดังนั้นจึงจำเป็นต้องนำผลรวมของค่าความรู้สึกแต่ละประเภทไปเป็นตัวหาร เพื่อให้ได้ค่าความน่าจะเป็นที่ถูกต้อง โดยการนำค่าที่ได้ของแต่ละประเภทความรู้สึกมาหาผลรวม ด้วยการทำซ้ำ (Foreach) ค่าคะแนนของแต่ละประเภทความรู้สึก (Score) ในแต่ละรอบค่า



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

ผลรวมของค่าคะแนนของแต่ละประเภทความรู้สึกจะถูกบวกเพิ่มด้วยค่าคะแนนของแต่ละประเภทความรู้สึกของคลาสนั้น จนกว่าจะครบทุกคลาส

ขั้นตอนที่ 4 หาค่าคะแนนความรู้สึก เป็นการหาค่าความน่าจะเป็นที่แท้จริงของโอกาสที่จะเกิดขึ้นของแต่ละประเภทความรู้สึก ด้วยการนำค่าคะแนนของแต่ละประเภทความรู้สึก (Score) มารวบรวมของค่าคะแนนของแต่ละประเภทความรู้สึก (Total Score) แล้วเก็บกลับคืนไว้ที่ ค่าคะแนนของแต่ละประเภทความรู้สึก (Score) จนกว่าจะครบทุกคลาส

ขั้นตอนที่ 5 เรียงลำดับและส่งออก นำค่าคะแนนของแต่ละประเภทความรู้สึก (Score) มาเรียงด้วยฟังก์ชัน SortDesc หลังจากที่เราเรียงเสร็จแล้ว นำค่าของผลการเรียงลำดับมาหาค่าที่มากที่สุดด้วยฟังก์ชัน Max หากค่าของคลาสใดมีคะแนนมากที่สุด หมายความว่าข้อความนี้มีค่าเป็นคลาสนั้นหรือข้อความนี้มีความรู้สึกตามประเภทที่มีค่ามากที่สุด และส่งออกผลการหาค่าคะแนนพร้อมระบุว่าข้อความประเภทใด จากขั้นตอนทั้งหมดที่กล่าวไว้ข้างต้นสามารถเขียนเป็นแผนภาพรหัสเทียม (Pseudo code) ได้ดังภาพที่ 5

```
Input: Words
Output: Sentiment Classification
Set Classes = {positive, negative, neutral}
Set Prior = {positive = 0.3333, negative = 0.3333, neutral = 0.3333}
Set Score = 1
Foreach(words as word)
    If (word == dictionary's word)
        Score[Class] += Score[Class]
    Endif
End foreach
Score[Class] = Prior[Class] * Score[Classes]
Foreach (Classes as Class)
    Totalscore += Score[class]
End foreach
Foreach (Classes as class)
    Score[Class] = Score[Class] / Totalscore
End foreach
Classification = Max(SortDesc(Score[Class]))
Return Classification
```

ภาพที่ 5 ขั้นตอนกระบวนการวิเคราะห์ความรู้สึกจากข้อความ



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8
The 8th STOU National Research Conference

จากภาพที่ 5 เป็นขั้นตอนกระบวนการวิเคราะห์ความรู้สึกจากข้อความซึ่งมีความสอดคล้องกับ 5 ขั้นตอนข้างต้น โดยนำเข้าข้อมูลที่ผ่านกระบวนการตัดคำมาแล้ว มาคำนวณหาค่าคะแนนและส่งออกผลลัพธ์การวิเคราะห์ความรู้สึกว่าข้อความเป็นประเภทใด

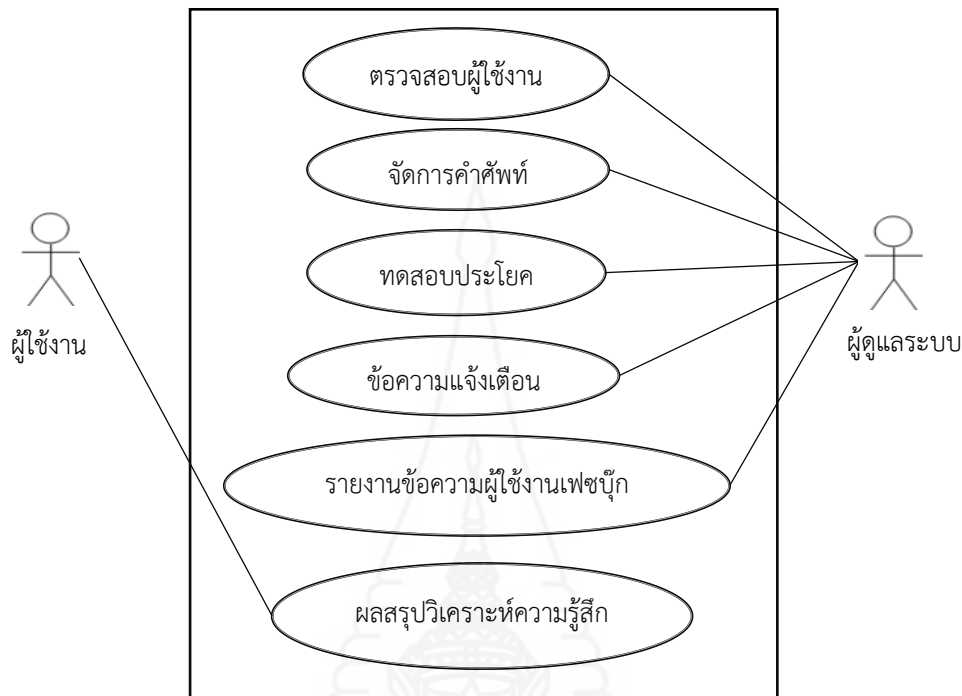
ตารางที่ 2 ผลการทดสอบการตัดข้อความ การหาค่าคะแนนแต่ละประเภทความรู้สึกและผลลัพธ์การวิเคราะห์ความรู้สึก

ตัวอย่างข้อความ	การตัดคำ	ค่าคะแนน (pos/neg/neu)	ผลลัพธ์
ภูมิใจมากที่ได้เรียนที่มสธ	ภูมิใจ มาก ที่ ได้ เรียน ที่ มสธ	0.5000/0.2500/0.2500	Positive
ดีใจที่ได้เป็นศิษย์มสธ	ดีใจ ที่ ได้ เป็น ศิษย์ มสธ	0.4167/0.3333/0.2500	Positive
แชร์ได้ไหมครับท่านอาจารย์	แชร์ ได้ ไหม ครับ ท่าน อาจารย์	0.2667/0.2667/0.4667	Neutral
เปิดรับสมัครวันไหนครับ	เปิด รับ สมัคร วัน ไหน ครับ	0.2500/0.2500/0.5000	Neutral
ยังไม่ได้รับหนังสือเลยคะ	ยัง ไม่ ได้ รับ หนังสือ เลย คะ	0.2500/0.5000/0.2500	Negative
ข้อสอบยากมาก	ข้อ สอบ ยาก มาก	0.2500/0.5000/0.2500	Negative

6. การออกแบบการพัฒนาระบบ โดยผู้วิจัยได้พัฒนาระบบในรูปแบบของเว็บไซต์ (Web site) ที่พัฒนาโดยภาษา PHP และใช้ MySQL เป็นระบบจัดการฐานข้อมูล โดยระบบที่พัฒนาขึ้นได้พัฒนาส่วนติดต่อผู้ใช้งานที่ต้องการวิเคราะห์ความรู้สึก การประมวลผลการวิเคราะห์ความรู้สึกจากข้อความแสดงความคิดเห็นภาษาไทย และส่วนแสดงผลการวิเคราะห์ความรู้สึกจากข้อความแสดงความคิดเห็นภาษาไทยที่แสดงเป็นเชิงบวกและเชิงลบ พร้อมทั้งการส่งข้อความแจ้งเตือนไปยังไลน์ (Line) ของผู้ดูแลระบบ โดยใช้ Line Notify ซึ่งเป็นบริการที่ทาง Line ได้เตรียมไว้ให้ในรูปแบบของ API (Application Programing Interface) ให้กับนักพัฒนาระบบสามารถนำไปใช้ต่อยอดพัฒนาระบบงาน โดยได้มีการออกแบบแผนภาพยูสเคส ไดอะแกรม (Use Case Diagram) (วฤชาชัย ร่มสายหยุด, 2561) เพื่อแสดงขอบเขตสิทธิ์การใช้งานระบบ ซึ่งประกอบด้วยผู้ใช้ระบบ (Actor) ดังนี้ คือ ผู้ดูแลระบบ (Admin) และผู้ใช้งาน (User) ดังแสดงในภาพที่ 6



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8
The 8th STOU National Research Conference



ภาพที่ 6 ยูสเคสไดอะแกรมการพัฒนาระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time)

6. การประเมินผล การประเมินผลและการทดสอบเพื่อประเมินประสิทธิภาพการจำแนกข้อมูลในงานวิจัยนี้ จะใช้การวัดค่าความแม่นยำ (Precision) ค่าเรียกคืน (Recall) ค่าความถูกต้อง (Accuracy) และการวัดประสิทธิภาพโดยรวม (F-measure) เพื่อทดสอบความถูกต้อง โดยเริ่มพิจารณาค่าความถูกต้องจากตาราง Confusion Matrix ซึ่งเป็นการประเมินผลลัพธ์จากระบบเปรียบเทียบกับผลลัพธ์จริง ๆ ที่ได้จากมนุษย์ ใช้วิธีคำนวณค่า Precision, Recall, Accuracy และ F-measure โดยมีวิธีการคำนวณค่าดังสมการที่ 1 ถึงสมการที่ 4

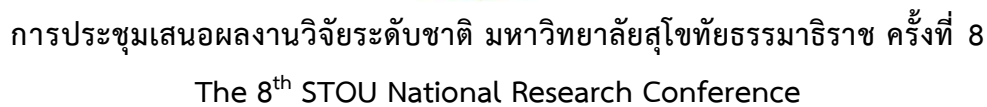
$$\text{Precision} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}} \quad (1)$$

$$\text{Recall} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Negative}} \quad (2)$$

$$\text{Accuracy} = \frac{\text{True Positive} + \text{True Negative}}{\text{Total Data Set}} \quad (3)$$

$$\text{F-measure} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

โดยที่ True Positive (TP) คือ จำนวนของความคิดเห็นที่เป็นเชิงบวกที่ระบบว่าเป็นเชิงบวก
True Negative (TN) คือ จำนวนความคิดเห็นที่เป็นเชิงลบที่ระบบว่าเป็นเชิงลบ
False Positive (FP) คือ จำนวนความคิดเห็นที่เป็นเชิงลบที่ระบบว่าเป็นเชิงบวก
False Negative (FN) คือ จำนวนความคิดเห็นเป็นเชิงบวกที่ระบบว่าเป็นเชิงลบ
Total Data Set คือ จำนวนของข้อมูลที่นำมาทดสอบทั้งหมด



ผู้วิจัยได้วิเคราะห์และออกแบบระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ของนักศึกษานพชนบักโดยใช้ตัวจำแนกข้อมูลนาอ์พ เบย์ สำหรับภาษาไทย โดยนำข้อความที่แสดงความคิดเห็นพชนบักพจนมหาวิทยาลัยสุโขทัยธรรมาธิราชมาวิเคราะห์ความรู้สึก (Sentiment Analysis) ในรูปแบบเว็บแอปพลิเคชัน (Web Application) โดยแบ่งการทำงานออกเป็น 2 ส่วนหลัก คือผู้ใช้งานระบบและพ้ดแลระบบ ซึ่งมีรายละเอียดดังนี้

1.1 ส่วนผู้ใช้งานระบบ สามารถดูผลที่ได้จากการวิเคราะห์ข้อมูลโดยแสดงผลในรูปแบบสถิติจำนวนของข้อความความคิดเห็น (COMMENTS) จำนวนของผู้กดชื่นชอบเพชบักเพจ (LIKES) จำนวนความเห็นเชิงบวก (GOOD) จำนวนความเห็นเชิงลบ (BAD) และส่วนแสดงรายละเอียดความคิดเห็น ซึ่งจะแสดงทีละ 20 ข้อความ โดยระบบจะแสดงผลลัพท์ด้วยวิธีดึงข้อมูลชุดใหม่ขึ้นมาแสดงทดแทนข้อมูลชุดเดิมแบบสลับ ส่วนข้อมูลชุดเดิมจะถูกซ่อนไปที่ละ 1 ข้อความแบบเรียลไทม์ (Real-time) โดยแบ่งเป็นความเห็นเชิงบวก 10 และความเห็นเชิงลบ 10 ข้อความ โดยแสดงผลดังภาพที่ 7

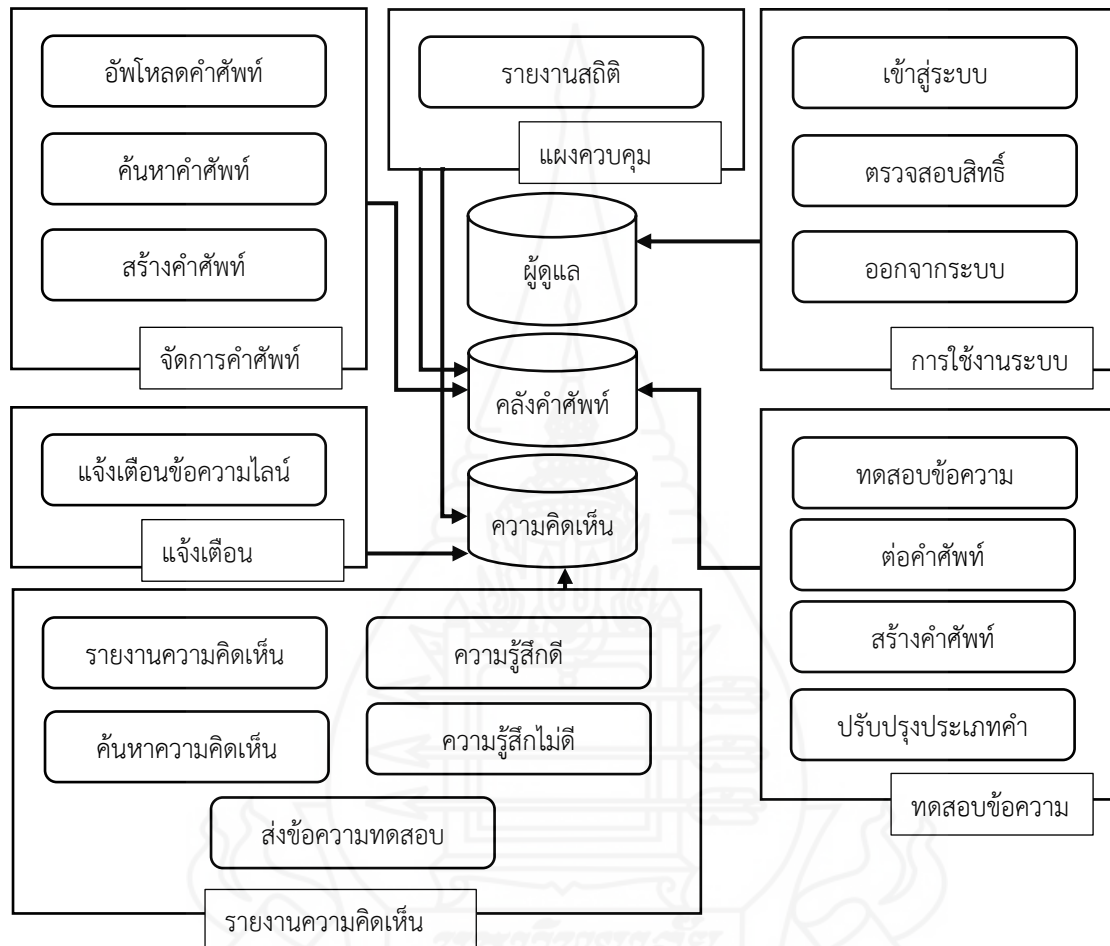
ภาพที่ 7 หน้าจอแสดงผลของระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time) ส่วนผู้ใช้งานระบบ

532



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8
The 8th STOU National Research Conference

ความรู้สึกเชิงลบหรือความรู้สึกที่เป็นกลาง รวมไปถึงการสมัครรับการแจ้งเตือนข้อความผ่านทางไลน์ (Line) โดยแบ่งออกเป็น 6 คุณลักษณะ (Feature) โดยแต่ละคุณลักษณะมีรายละเอียดดังภาพที่ 8



ภาพที่ 8 ภาพรวมของส่วนผู้ดูแลระบบ

2. ผลการประเมินประสิทธิภาพของระบบ

งานวิจัยมีการทดสอบประสิทธิภาพของระบบด้วยค่าความแม่นยำ (Precision) ค่าเรียกคืน (Recall) ค่าความถูกต้อง (Accuracy) และการวัดประสิทธิภาพโดยรวม (F-measure) เพื่อทดสอบความถูกต้องของข้อมูล ผู้วิจัยเลือกกระบวนการทดสอบโดยใช้โปรแกรม Rapid Miner ในการช่วยทดสอบ เปรียบเทียบการทำงานระหว่างผลการวิเคราะห์ของ ผู้พัฒนาระบบและผลการวิเคราะห์จากระบบ กระบวนการที่ทำให้ได้ค่าความถูกต้องที่มากนั้น ผู้วิจัยได้ใช้วิธีทดสอบที่ละ ประโยคด้วยตัวเองผ่านทางเครื่องมือที่พัฒนาขึ้นมา ตัวแปรสำคัญที่ทำให้เกิดข้อผิดพลาดมากที่สุดคือ คำศัพท์ที่ระบบไม่รู้จำ ทำให้ไม่สามารถให้ค่าของคำนั้นได้ ผู้วิจัยจึงได้ออกแบบเครื่องมือเพิ่มคำศัพท์เพื่อให้ระบบได้เรียนรู้ หลักจากที่เพิ่มคำศัพท์ใหม่ เข้าไปในระบบแล้ว จากนั้นก็ทำการทดสอบอีกครั้ง จนกว่าระบบจะทำนายผลออกมาได้ค่าที่ดีที่สุด มีบางประโยคที่ระบบ



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8

The 8th STOU National Research Conference

มักจะทำนายผลผิดพลาดคือประโยคในลักษณะมีความขัดแย้งกันเองในรูปประโยค ซึ่งผลลัพธ์ที่ระบบทำนายออกมามักจะเป็นค่าเชิงบวก เพราะค่าลำดับแรกของอาร์เรย์ (Array) ของคลาส (Class) คือค่าเชิงบวก (Positive) ผู้วิจัยใช้ข้อมูลทดลองจำนวน 12,548 ข้อความ โดยการเปรียบเทียบค่าการวิเคราะห์จากระบบและส่วนที่วิเคราะห์จากผู้พัฒนาระบบทั้ง 12,548 ข้อความ หลังจากที่ได้ผลการวิเคราะห์ทั้งสองส่วนแล้ว ผู้วิจัยนำข้อมูลที่ได้เข้าไปทดสอบด้วยโปรแกรม Rapid miner โดยให้ผลการวิเคราะห์จากผู้พัฒนาระบบเป็นข้อมูล Training Data ส่วนผลการวิเคราะห์ด้วยระบบเป็นข้อมูลชุดทดสอบ (Testing Data) ผลการทดสอบได้ผลลัพธ์โดยแยกออกเป็นผลการวิเคราะห์เชิงบวก (Positive) 97.75% ผลการวิเคราะห์เชิงลบ (Negative) 93.77% ผลการวิเคราะห์เป็นกลาง (Neutral) 97.96% และผลการทดสอบค่าความถูกต้อง 97.60% ดังตารางที่ 3

ตารางที่ 3 ผลการทดสอบประสิทธิภาพด้วยโปรแกรม Rapid Miner

Accuracy (97.60%)	True Positive	True Neutral	True Negative	Class precision
Pred. Positive	5,725	89	34	97.90%
Pred. Neutral	103	4,953	11	97.76%
Pred. Negative	28	14	677	94.16%
Class Recall	97.76 %	97.96%	93.77 %	

อภิปรายผลการวิจัย

บทสรุปงานวิจัยนี้ผู้วิจัยมุ่งเน้นการนำเสนอการวิเคราะห์ออกแบบและพัฒนาระบบวิเคราะห์ความรู้สึกแบบเรียลไทม์ (Real-time) ของนักศึกษาบนเฟซบุ๊กโดยใช้ตัวจำแนกข้อมูลนาอ์ฟ เบย์ สำหรับภาษาไทย และประเมินค่าความแม่นยำและค่าความถูกต้องของระบบวิเคราะห์ความรู้สึก สามารถสรุปผลการวิจัยซึ่งเป็นผลทดสอบประสิทธิภาพของระบบ เพื่อทดสอบความถูกต้อง โดยแบ่งออกเป็นผลวิเคราะห์ความคิดเห็นเชิงบวก (Positive) 97.76% , ผลวิเคราะห์ความคิดเห็นเชิงลบ (Negative) 93.77% , ผลวิเคราะห์ความคิดเห็นเป็นกลาง (Neutral) 97.96% , ค่าความแม่นยำ (Precision) 96.61% , ค่าเรียกคืน (Recall) 96.50 % , ค่าความถูกต้อง (Accuracy) 97.60% และค่าผลการทดสอบภาพรวมของระบบ (F-measure) 96.55% สำหรับส่วนที่ทำให้ระบบทำนายผลผิดพลาดนั้นเกิดจากสาเหตุส่วนใหญ่คือ ประโยคที่มีลักษณะขัดแย้งกันเช่น “ได้รับเร็วขึ้นสำหรับบางคน ส่วนของเราช้าเหมือนเดิม” ระบบทำนายผลออกมาเป็นเชิงบวกแต่ที่ถูกต้องควรเป็นเชิงลบ และอีกปัญหาหนึ่งคือผู้ใช้งานเฟซบุ๊กพิมพ์ข้อความผิดทำให้ระบบไม่สามารถแปลความหมายที่ถูกต้องได้ จากกระบวนการพัฒนาระบบที่นำเสนอยังสามารถนำไปประยุกต์ใช้เป็นเครื่องมือสำรวจความพึงพอใจและติดตามทัศนคติของสาธารณะที่มีต่อผลิตภัณฑ์และบริการ เพื่อเพิ่มขีดความสามารถทางการแข่งขันให้ธุรกิจและหน่วยงานรัฐ ระบบสามารถวิเคราะห์ข้อความที่ใช้ภาษาพูดที่ไม่เป็นทางการได้อย่างมีประสิทธิภาพ สามารถติดตามความเคลื่อนไหว หรือการตอบรับของกลุ่มเป้าหมาย พร้อมทั้งระบบแจ้งเตือนผ่านข้อความไลน์ (Line) เพื่อให้สามารถตอบสนองความต้องการที่แท้จริง และพร้อมรับมือกับปัญหาได้อย่างทันทั่วทั้ง



การประชุมเสนอผลงานวิจัยระดับชาติ มหาวิทยาลัยสุโขทัยธรรมาธิราช ครั้งที่ 8 The 8th STOU National Research Conference

ข้อเสนอแนะ

1. พจนานุกรมการแบ่งสถานะของคำยังมีน้อยสำหรับการนำมาใช้ในการวิเคราะห์ข้อมูลจำนวนมากดังนั้นก็ควรมีการเพิ่มคำพร้อมสถานะของคำให้มากขึ้นจากเดิมที่มีเพียง 19,080 คำ เพื่อให้ระบบสามารถนำมาใช้ในการวิเคราะห์ได้อย่างมีประสิทธิภาพมากขึ้น
2. ควรมีระบบแก้ไขคำผิดแบบอัตโนมัติเพราะเมื่อมีคำผิดจะทำให้ระบบไม่สามารถนำไปวิเคราะห์ได้เพราะระบบจะตัดคำนั้นออกเป็นคำที่ไม่รู้จัก ทำให้การทำงานของระบบประสิทธิภาพลดลง
3. เฟซบุ๊กมีการปรับนโยบายการเข้าใช้งานข้อมูลผ่าน Facebook Graph API ซึ่งหากจะนำระบบไปใช้งานงานต่อจำเป็นต้องขออนุญาตการใช้งานข้อมูลไปทางเฟซบุ๊กหรือ เป็นผู้ดูแลเฟซบุ๊กเพจจะได้สิทธิ์ในการใช้งานข้อมูลโดยอัตโนมัติ

เอกสารอ้างอิง

- นภดล สิทธิเลิศ. (2556). การพยากรณ์อารมณ์ความรู้สึกจากการโพสต์ข้อความบนเฟซบุ๊ก (Facebook). สารนิพนธ์วิทยาศาสตรมหาบัณฑิต สาขาวิชาวิทยาการคอมพิวเตอร์ ภาควิชาวิทยาการคอมพิวเตอร์และสารสนเทศ บัณฑิตวิทยาลัย มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ.
- นริศร์ พรหมบุตร. (2550). การทำเหมืองข้อความคิดเห็นสินค้า กรณีศึกษาโทรศัพท์มือถือ. สารนิพนธ์วิทยาศาสตรมหาบัณฑิต สาขาวิชาวิทยาการคอมพิวเตอร์ ภาควิชาวิทยาการคอมพิวเตอร์และสารสนเทศ คณะวิทยาศาสตร์ประยุกต์ บัณฑิตวิทยาลัย สถาบันเทคโนโลยีพระจอมเกล้าพระนครเหนือ.
- พนิดา ทรงรัมย์. (2559). การจำแนกความคิดเห็นทางการเมืองบนเครือข่ายสังคมออนไลน์ โดยใช้วิธีการจำแนกแบบสัมพันธ์. วารสารวิทยาศาสตร์และเทคโนโลยี มทร.ธัญบุรี. Vol.6 No.1: 83-93.
- วฤชาย รมสายหยุด. (2561). หลักการวิศวกรรมซอฟต์แวร์. นนทบุรี : มหาวิทยาลัยสุโขทัยธรรมาธิราช.
- wp (2561). สถิติผู้ใช้ดิจิทัลทั่วโลก “ไทย” เติบโตขึ้นมากสุดในโลก-“กรุงเทพ” เมืองผู้ใช้ Facebook สูงสุด คำนวณวันที่ 7 กรกฎาคม 2561 จาก <https://www.brandbuffet.in.th/2018/02/global-and-thailand-digital-report-2018/>
- Daniel Jurafsky and James H. Martin (2018). Naive Bayes and Sentiment Classification, Speech and Language Processing. 61-79