

Thai-textual Cyberbullying Detection using Support Vector Machines

การพัฒนาแบบจำลองในการตรวจจับข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์ โดยใช้วิธีซัพพอร์ตเวกเตอร์แมชชีน

Received 18 Sep 19
Reviewed 27 Sep 19
Revised 25 Oct 19
Accepted 30 Oct 19

Nuttakit Jenkarn and Mahasak Ketcham*

ณัฐกิจ เจนการ¹ และ มหศักดิ์ เกตุฉ่ำ^{2*}

¹Guru Services Company Limited

¹บริษัท กูรูเซอร์วิส จำกัด

²Department of Information Technology, Faculty of Information Technology, King Mongkut's University of Technology North Bangkok, Thailand

²ภาควิชาสาขาการจัดการเทคโนโลยีสารสนเทศ คณะเทคโนโลยีสารสนเทศ มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ กรุงเทพมหานคร

*Corresponding Author, Tel. -, E-mail: mahasak.k@it.kmutnb.ac.th

*ผู้พิมพ์ประสานงาน โทรศัพท์ - อีเมล: mahasak.k@it.kmutnb.ac.th

Abstract

This research applied data mining and machine learning to develop a prototype “Thai-textual Cyberbullying Detection using Support Vector Machines” that can be future online Cyberbullying social media detection. Support Vector Machines (SVMs), K-Nearest Neighbour and Naïve Bayes are selected to use in our research. From the evaluation by Confusion Matrix, SVMs had the highest classification accuracy as 83.91

Keywords: Text Mining, Cyberbullying, Machine learning

บทคัดย่อ

งานวิจัยฉบับนี้ได้นำเอาเทคนิควิธีการทำเหมืองข้อความและการเรียนรู้ของเครื่อง มาประยุกต์ใช้ในการพัฒนาแบบจำลองจำแนกข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์ เพื่อให้สามารถนำไปพัฒนาต่อยอดในการป้องกันการกลั่นแกล้งทางไซเบอร์บนสื่อโซเชียลออนไลน์ได้ ซึ่งในงานวิจัยชิ้นนี้ได้ทำการพัฒนาแบบจำลองโดยใช้อัลกอริทึม Support Vector Machine , K-Nearest Neighbour และอัลกอริทึม Naïve Bayes โดยได้มีการประเมินประสิทธิภาพและทำการเปรียบเทียบประสิทธิภาพของอัลกอริทึม โดยใช้ Confusion Matrix โดยผลที่ได้คืออัลกอริทึม Support Vector Machine มีประสิทธิภาพในการจำแนกสูงที่สุด ซึ่งมีค่าความถูกต้องอยู่ที่ร้อยละ 83.91

คำสำคัญ: วิธีการทำเหมืองข้อความ การกลั่นแกล้งทางไซเบอร์ การเรียนรู้ของเครื่อง

1. บทนำ

ปัจจุบันนี้ทุกคนต่างก็มีบัญชีโซเชียลต่าง ๆ อย่าง Facebook, Instagram, Youtube หรือ Twitter ซึ่งผู้คนส่วนใหญ่นั้น เลือกที่จะรับข่าวสารต่าง ๆ จากทางสื่อ

โซเชียลมีเดียเหล่านี้ เพราะมีความสะดวกรวดเร็ว และเข้าถึงง่าย และยังสามารถเผยแพร่ต่อได้ง่ายอีกด้วย และด้วยการแพร่กระจายที่ง่ายและรวดเร็วของข่าวต่าง ๆ นี้เอง ทำให้เกิดการเผยแพร่ทั้งข่าวที่เป็นความจริง และไม่เป็น

ความจริง รวมถึงการแสดงความคิดเห็นต่างที่มีต่อข่าวเหล่านั้นทั้งในแง่บวกและแง่ลบนั้นมีมากขึ้น และด้วยเหตุนี้เองก็ทำให้มีผู้ที่ได้รับผลกระทบมากมายจากสื่อโซเชียลมีเดียเหล่านี้ WEF Global press release ได้เผยแพร่ผลงานวิจัยระดับโลกในเรื่องพลเมืองดิจิทัลของโลกนี้ [1] โดยระบุว่า เด็กไทยมีโอกาสเสี่ยงภัยจากออนไลน์ถึง (60%) ขณะที่ค่าเฉลี่ยทั่วโลกอยู่ที่ (56%) ซึ่งภัยออนไลน์ของเด็กไทยที่พบมากที่สุดมี 4 ประเภท คือ 1.Cyber bullying (49%), 2. การเข้าถึงสื่อลามกและพูดคุยเรื่องเพศกับคนแปลกหน้าในโลกออนไลน์ (19%), 3. ติดเกม (12%) และ 4.ถูกล่อลวงออกไปพบคนแปลกหน้า (7%) จะเห็นได้ว่า Cyberbullying ในเด็กไทยนั้นเป็นเรื่องที่น่าเป็นห่วง และยังมีค่าที่สูงกว่าค่าเฉลี่ยโลกที่อยู่ (47%) อีกด้วย

Cyberbullying [2] คือ การรังแกผู้อื่นผ่านทางโลกออนไลน์ ซึ่งการรังแกนั้นเป็นไปทั้งในรูปแบบของการคำพูด, การกล่าวหา, การใช้ถ้อยคำเสียชื่อเสียงต่อบุคคลเป้าหมาย และมีความโน้มในการรังแกที่ต่อเนื่องมากกว่าหนึ่งครั้ง ซึ่งการกระทำนี้สามารถเกิดขึ้นได้บ่อยกว่าการแก่งัดกันทั่วไป เพราะในโลกออนไลน์นั้น ผู้รังแกไม่ได้เผชิญหน้ากับเป้าหมายจริง ๆ โดยผลจากการวิจัย เรื่อง “ความชุกและปัจจัยที่เกี่ยวข้องกับการกลั่นแกล้งบนโลกโซเชียล ในระดับชั้น ม.1-3” โดยเป็นความร่วมมือในระดับนานาชาติ 14 ประเทศทั่วโลก [3] ชี้ให้เห็นว่า เด็กไทยนั้นมีประสบการณ์ถูกกลั่นแกล้งในชีวิตจริงมากถึง 80 % ซึ่งมากกว่าประเทศอย่างสหรัฐอเมริกา ยุโรป และญี่ปุ่นถึง 4 เท่า โดยพบารูปแบบในการรังแกที่พบว่าถูกกระทำมากที่สุด คือ การโดนล้อเลียนและการถูกตั้งฉายาที่ 79.4% การถูกเพิกเฉย ไม่สนใจ 54.4%และคนอื่นไม่เคารพ 46.8% ตามด้วยการถูกปล่อยข่าวลือ การถูกนำรูปไปตัดต่อ ถูกข่มขู่ และการถูกทำให้หวาดกลัว ตามลำดับ โดยส่วนใหญ่มากถึง 89.2% เลือกที่จะปรึกษาเพื่อนมากกว่าคนในครอบครัว และด้วยความตระหนักถึงปัญหาเรื่องการ Cyberbullying ของทั่วโลก จึงได้มีการกำหนดให้วันศุกร์ที่ 3 ของเดือนมิถุนายนเป็นวันหยุดการกลั่นแกล้งบนโลกออนไลน์ (Stop Cyberbullying Day) [4] เพื่อทำให้สังคมได้ตระหนักถึงปัญหาเหล่านี้

ผลการวิจัยเรื่อง Cyberbullying classification using text mining [5] ได้มีการนำเสนอการสร้างแบบจำลองในการจำแนกประเภทของการกลั่นแกล้งทางอินเทอร์เน็ตในการสนทนาด้วยข้อความ(ภาษาอังกฤษ)โดยใช้วิธีการทำ

เหมือนข้อความ ซึ่งได้เสนออัลกอริทึมที่ใช้คือ Support Vector Machine เปรียบเทียบกับ Naive Bayes ซึ่งผลการวิจัยชี้ให้เห็นว่าวิธี Support Vector Machine นั้นให้ค่าความแม่นยำที่สูงกว่าวิธี Naive Bayes และมีความสามารถที่จะจำแนกประเภทของการกลั่นแกล้งทางอินเทอร์เน็ตในการสนทนาด้วยข้อความ(ภาษาอังกฤษ) ได้เป็นอย่างดี และอีกหนึ่งงานวิจัยของไทยเรื่องการวิเคราะห์ความคิดเห็นภาษาไทยเกี่ยวกับการรีวิวสินค้าออนไลน์ โดยใช้ขั้นตอนวิธีซัพพอร์ตเวกเตอร์แมชชีน [6] โดยได้นำเสนอการเปรียบเทียบประสิทธิภาพการวิเคราะห์ความคิดเห็นภาษาไทย ของอัลกอริทึม 4 ตัว ได้แก่ ซัพพอร์ตเวกเตอร์แมชชีน ต้นไม้ตัดสินใจ นาอีฟเบย์ และเคเนียร์เซนเนอร์ ซึ่งจากผลการวิจัยพบว่าอัลกอริทึมซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine) ให้ประสิทธิภาพในการจำแนก คุณลักษณะที่ดีที่สุด ดังนั้นงานวิจัยในครั้งนี้จึงมีแนวคิดที่จะใช้เทคนิคการสร้างแบบจำลองในการจำแนก

ข้อความโดยใช้วิธีซัพพอร์ตเวกเตอร์แมชชีน มาใช้ในการพัฒนาแบบจำลองในการตรวจจับข้อความภาษาไทยที่เป็น การกลั่นแกล้งทางโซเชียลมีเดีย โดยใช้วิธีซัพพอร์ตเวกเตอร์แมชชีน

2. วัสดุ อุปกรณ์ และวิธีการวิจัย

2.1 การรวบรวมข้อมูล

ผู้วิจัยได้มีการเก็บรวบรวมข้อมูล ซึ่งเป็นข้อความแสดงความคิดเห็นของผู้ใช้สื่อสังคมออนไลน์ทั้ง Facebook , YouTube , Twitter และ Instagram ที่เป็นการแสดงความเห็นต่อบุคคลที่เป็นกระแสมือในขณะนั้น เช่น ดารา นักกีฬา หรือ ผู้ต้องหาคดีต่าง ๆ เป็นต้น โดยใช้เครื่องมือ (API) ของสื่อสังคมออนไลน์นั้น ๆ ในการดึงข้อความแสดงความคิดเห็นและนำมาจัดเก็บให้อยู่ในรูปแบบของไฟล์ซีเอสวี (CSV)

ตารางที่ 1 ตัวอย่างข้อมูลที่รวบรวมได้

Text
สุดยอด...ของความตอแหลจังไร
ดูไปเขินไป เริ่มจะรักผู้ชายคนนี้แล้วเนี่ย
ตอแหล...หน้าด้านไม่เบาอย่างหนัาระลั่นได้เนมิง
อิตอแหล!!! ออกไปจากวงการอิตอก...นีสัยแย
เป็นผู้ชายที่พูดครบได้ดูน่าเอ็นดูมากครับ

2.2 การจัดรูปแบบข้อมูล

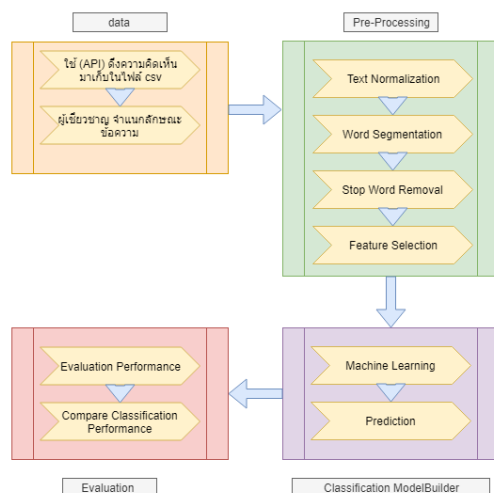
ในการจำแนกข้อความแสดงความคิดเห็นภาษาไทยที่เป็นการกลั่นแกล้งทางโซเชียลนั้น ผู้วิจัยได้แบ่งการจำแนกออกเป็น 2 ลักษณะ ได้แก่ ลักษณะข้อความที่เป็นการกลั่นแกล้งทางโซเชียล และลักษณะข้อความที่ไม่เป็นการกลั่นแกล้งทางโซเชียล ซึ่งผู้วิจัยได้นำข้อความแสดงความคิดเห็นภาษาไทยที่สามารถเก็บรวบรวมได้ มาวิเคราะห์ตามลักษณะที่ได้จำแนกขึ้น ซึ่งจะวิเคราะห์ทีละข้อความว่าจัดอยู่ในลักษณะใด โดยให้ผู้เชี่ยวชาญเป็นผู้ทำการวิเคราะห์และบันทึกเป็นตารางบนไฟล์ซีเอสวี (CSV)

2.3 การออกแบบวิธีการวิเคราะห์

ตารางที่ 2 การกำหนด Class ให้กับข้อความ

Text	Class
สุดยอด...ของความตอแหลจังไร	Yes
ดูไปเงินไป เริ่มจะรักผู้ชายคนนี้แล้วเนี่ย	No
ตอแหล...หน้าด้านไม่เบายังหน้าระลั่นได้เนะมึง	Yes
อืตอแหล!!! ออกไปจากวงการอืตอก...นี้เสียแย	Yes
เป็นผู้ชายที่พูดครบได้ดูน่าเอ็นดูมากครับ	No

กระบวนการการพัฒนาแบบจำลองในการตรวจจับข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางโซเชียลและการเปรียบเทียบอัลกอริทึมที่ได้เลือกมาทดสอบ สามารถอธิบายขั้นตอนได้ดังนี้



รูปที่ 1 กระบวนการการพัฒนาแบบจำลองในการตรวจจับข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางโซเชียล

2.3.1 ขั้นตอนการเตรียมข้อมูล (Data) ซึ่งประกอบไปด้วย

2.3.1.1 เก็บรวบรวมข้อความแสดงความคิดเห็นจากการวิเคราะห์ลักษณะจากสื่อสังคมออนไลน์ โดยใช้เครื่องมือ (API) ในการดึงข้อความแสดงความคิดเห็นจาก YouTube และ Twitter และใช้ในการดึงข้อมูลด้วยตนเองจาก Facebook และ Instagram แล้วจัดเก็บในไฟล์ CSV

2.3.1.2 ให้ผู้เชี่ยวชาญตรวจสอบและวิเคราะห์ข้อความแสดงความคิดเห็นที่ได้จาก Facebook, YouTube, Twitter และ Instagram และจำแนกลักษณะข้อความว่าเป็นลักษณะข้อความที่เป็นการกลั่นแกล้งทางโซเชียล หรือลักษณะข้อความที่ไม่เป็นการกลั่นแกล้งทางโซเชียล แล้วจัดเก็บในไฟล์ CSV

2.3.2 นำข้อมูลเข้าสู่กระบวนการก่อนหน้า (Pre-Processing) ซึ่งประกอบไปด้วย

2.3.2.1 Text Normalization ทำการลบสัญลักษณ์พิเศษออกจากข้อความ ซึ่งในที่นี้จะรวมถึงสัญลักษณ์ของภาษาไทยและตัวเลขต่างที่ไม่ต้องการด้วยเช่น @, #, \$, %, ๖, ๕ เป็นต้นรวมทั้งช่องว่าง เพื่อกำจัดข้อมูลส่วนเกินออก

2.3.2.2 การตัดคำ (Word Segmentation) การตัดคำนั้นเป็นกระบวนการหลักและมีความสำคัญสำหรับภาษาไทย ซึ่งจะทำให้การตัดคำให้เป็นคุณลักษณะแบบคำเดี่ยว (Single word) ด้วยวิธีการตัดคำแบบสอดคล้องมากที่สุด (Maximum Matching algorithm) โดยใช้พจนานุกรม (Thai national corpus) เป็นตัวเปรียบเทียบโดยกำหนดให้ตัวอักษรของการตัดคำที่ตัดได้แยกจากกันด้วยสัญลักษณ์ไปป์ (|) แล้วนำไปบันทึกลงในไฟล์ข้อมูล CSV

ตารางที่ 3 การตัดคำเพื่อใช้ในการจำแนกความคิดเห็น

Text	Class
สุด ยอด ของ ความ ตอ แหล จัง ไร	Yes
ดู ไป เงิน ไป เริ่ม จะ รัก ผู้ ชาย คน นี้ แล้ว เน ีย	No
ตอ แหล หน้า ด้าน ไม่ เบา ยัง หน้า ระ ลั่น ได้ เน ะ มึง	Yes
อื ตอ แหล ออก ไป จาก วง การ อื ตอก นี้ เสีย แย	Yes
เป็น ผู้ ชาย ที่ พูด ครบ ได้ ดู น่า เอ็น ดู มาก ครับ	No

2.3.2.3 การกำจัดคำหยุด (Stop Word Removal) เป็นการนำคำที่ไม่มีความหมายหรือคำที่ไม่มีความสำคัญในเอกสารออก โดยที่ความหมายของคำหรือข้อความจะไม่เปลี่ยนแปลง คำหยุดจะปรากฏในข้อความทุกข้อความ ซึ่งคำหยุดเป็นคุณลักษณะที่ไม่มีความเกี่ยวข้องหรือไม่มีความ

ประโยชน์ ในการจำแนกหมวดหมู่ข้อความ ดังนั้นการจัดคำหุตุจึงเป็นกระบวนการที่ควรทำก่อนการจัดทำดัชนีเพื่อกำจัดคุณลักษณะที่ไม่เป็นประโยชน์ และลดขนาดของดัชนีลง ซึ่งจะช่วยประหยัดทั้งพื้นที่และเวลาในการประมวลผล

2.3.2.4 ขั้นตอนการเลือกคุณลักษณะ (Feature Selection)

คุณลักษณะจำนวนมากมาใช้ในการประมวลผลจะใช้เวลาและทรัพยากรในการประมวลผลมาก ดังนั้นจึงจำเป็นต้องมีการเลือกคุณลักษณะเพื่อให้ข้อมูลมีขนาดลดลง แต่ต้องสูญเสียลักษณะสำคัญของข้อมูลและความถูกต้องของผลลัพธ์ให้น้อยที่สุด สำหรับเทคนิคที่นำมาใช้ในการเลือกคุณลักษณะมีหลายเทคนิค ในงานวิจัยนี้ประยุกต์เทคนิคการเลือกคุณลักษณะด้วยวิธีการทางสถิติมาใช้ โดยใช้วิธีการให้น้ำหนักของคำในเอกสารแบบ TF/IDF (Term Frequency/Inverse Document Frequency) ซึ่งเป็นวิธีที่ใช้ประเมินความสำคัญของคำหรือคุณลักษณะนั้น ๆ ในเอกสารที่เป็นชุดฝึก (Training Set) ทั้งหมด (Jing, Huang and Shi, 2002) คุณลักษณะที่มีค่าน้ำหนัก TF/IDF มากหมายความว่า คุณลักษณะนั้นมีประสิทธิภาพในการจำแนกเอกสารมากเช่นกัน การคำนวณค่าน้ำหนัก TF-IDF แสดงในสมการที่ 1 และ 2

$$W_{(f,d)} = TF_{(f,d)} \times IDF_{(f)} \quad (1)$$

$$IDF_{(f)} = \log \frac{|D|}{|DF_{(f)}|} \quad (2)$$

โดยที่ $W_{(f,d)}$ คือ ค่าน้ำหนักของคุณลักษณะ (f) ที่ปรากฏในเอกสาร (d)

$TF_{(f,d)}$ คือ ความถี่ของคุณลักษณะ (f) ที่ปรากฏในเอกสาร (d)

$|D|$ คือ จำนวนเอกสารทั้งหมดในข้อมูลชุดฝึก (Training Set)

$|DF_{(f)}|$ คือ จำนวนของเอกสารทั้งหมดที่มีคุณลักษณะ (f) ปรากฏอยู่

การพิจารณาเลือกคุณลักษณะจะกระทำภายหลังจากคำนวณน้ำหนัก TF-IDF เสร็จ โดยพิจารณาจากคุณลักษณะที่ให้ค่า TF-IDF สูงในแต่ละ Class โดยนำมารวมกันและหมวดค่าที่ซ้ำกันเข้าด้วยกัน จากนั้นนำคุณลักษณะที่เลือกมาแปลงให้อยู่ในรูปแบบของเมตริกซ์เอกสาร (Document-

Terms Matrix) ดังตารางที่ 4 ก่อนนำไปสร้างแบบจำลองด้วยเทคนิคการเรียนรู้ของเครื่อง

ตารางที่ 4 เมตริกซ์เอกสาร

Row	ทำ	ตัว	...	เลว	Class
1	0.212	0	...	0.453	Yes
2	0.122	0.422	...	0	No
3	0	0.312	...	0.453	Yes
...	...	...	...	...	...
n	0.233	0	...	0	No

2.3.3 การสร้างแบบจำลองในการจำแนกข้อความ (Classification Modelbuilder)

ใช้อัลกอริธึม Support Vector Machine (SVM), K-Nearest Neighbour และ Naive Bayes ในการจำแนกข้อความออกเป็น 2 ประเภทคือ ข้อความที่เป็นการกลั่นแกล้งทางไซเบอร์(Yes) และข้อความที่ไม่เป็นการกลั่นแกล้งทางไซเบอร์ (No)

2.3.4 ขั้นตอนการประเมินประสิทธิภาพ (Evaluation)

2.3.4.1 ทำการประเมินประสิทธิภาพของตัวจำแนกแต่ละตัวที่ใช้ อัลกอริธึม Support Vector Machine (SVM), K-Nearest Neighbour และ Naive Bayes โดยใช้วิธี Cross-validation Test เพื่อหาค่าความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่าความระลึก (Recall) และค่าความถ่วงดุล (F-Measure) ของแต่ละตัวจำแนก

2.3.4.2 ทำการเปรียบเทียบประสิทธิภาพของตัวจำแนกแต่ละตัวโดยใช้ค่าความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่าความระลึก (Recall) และค่าความถ่วงดุล (F-Measure) ของแต่ละตัวจำแนกเป็นตัวเปรียบเทียบ

2.4 การประเมินประสิทธิภาพ

การประเมินประสิทธิภาพของแบบจำลอง จะใช้ Cross-validation Test โดยทำการแบ่งข้อมูลออกเป็นหลายส่วน โดยที่แต่ละส่วนมีจำนวนข้อมูลเท่ากัน แล้วนำข้อมูลหนึ่งส่วนไว้ใช้เป็นตัวทดสอบประสิทธิภาพของโมเดล และนำข้อมูลส่วนที่เหลือไปสร้างโมเดล ทำวนไปเช่นนี้จนครบจำนวนที่แบ่งไว้ แล้วคำนวณหาค่า Confusion Matrix ค่า

ความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่าความระลึก (Recall) และค่าความถ่วงดุล (F-Measure)

2.4.1 ค่า Confusion Matrix

2.4.1.1 True Positive (TP) หมายถึง จำนวนที่ทำนายตรงกับข้อมูลจริงในคลาสที่กำลังพิจารณา สมการ True Positive rate = $\sum \text{True positive} / \sum \text{Condition positive}$

2.4.1.2 True Negative (TN) หมายถึง จำนวนที่ทำนายตรงกับข้อมูลจริงในคลาสที่ไม่ได้กำลังพิจารณา สมการ True Negative rate = $\sum \text{True negative} / \sum \text{Condition negative}$

2.4.1.3 False Positive (FP) หมายถึง จำนวนที่ทำนายผิดเป็นคลาสที่กำลังพิจารณา สมการ False Positive rate = $\sum \text{False positive} / \sum \text{Condition negative}$

2.4.1.4 False Negative (FN) หมายถึง จำนวนที่ทำนายผิดเป็นคลาสที่ไม่ได้กำลังพิจารณา สมการ False Negative rate = $\sum \text{False negative} / \sum \text{Condition positive}$

2.4.2 ค่า Precision แสดง จำนวนที่ทำนายถูกจากข้อมูลที่ทำนายว่าเป็นคลาสที่กำลังพิจารณา สมการ Precision = $TP / (TP + FP)$

2.4.3 ค่า Re-call แสดง จำนวนข้อมูลที่ทำนายถูก สมการ Recall = $TP / (TP + FN)$

2.4.4 ค่า F-measure คือ ค่าเฉลี่ยของ Precision และ Recall สมการ F-measure = $[2(Precision \times Re-call) / (Precision + Re-call)]$

2.4.5 ค่า Accuracy คือ จำนวนข้อมูลที่ทำนายถูกของทุกคลาส สมการ Accuracy = $(TP + TN) / (TP + TN + FP + FN)$

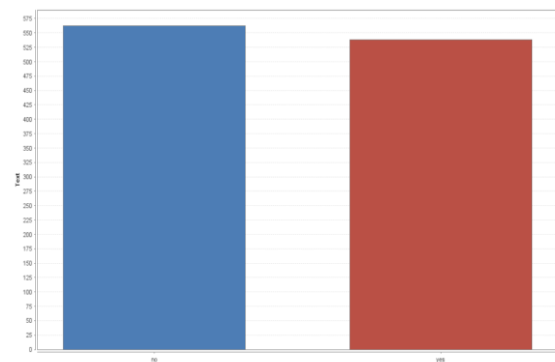
2.5 การเปรียบเทียบประสิทธิภาพ

ทำการเปรียบเทียบประสิทธิภาพของแบบจำลองในการตรวจจับข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์ระหว่างวิธี Support Vector Machine (SVM) วิธี K-Nearest Neighbour และ วิธี Naive Bayes โดยเปรียบเทียบค่าความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่าความระลึก (Recall) และค่าความถ่วงดุล (F-Measure) ที่วัดได้จากข้างต้นของทั้งสองวิธี เพื่อหาวิธีที่ดีที่สุด

3. ผลการวิจัย

3.1 การกำหนดกลุ่มของข้อความ

งานวิจัยได้รวบรวมข้อความจากสื่อสังคมออนไลน์และได้ทำการแบ่งกลุ่มข้อความออกเป็น 2 กลุ่ม คือข้อความที่เป็นการกลั่นแกล้ง (yes) และข้อความที่ไม่เป็นการกลั่นแกล้ง (no) โดยมีกลุ่มตัวอย่างทั้งหมด 1,100 ข้อความ โดยผู้เชี่ยวชาญในการวิเคราะห์และกำหนดกลุ่มของข้อความซึ่งสามารถกำหนดได้ดังนี้ ข้อความที่เป็นการกลั่นแกล้ง (yes) 538 ข้อความ และข้อความที่ไม่เป็นการกลั่นแกล้ง (no) 562 ข้อความ



รูปที่ 2 กลุ่มข้อความจากสื่อสังคมออนไลน์

3.2 ผลการประเมินการเรียนรู้ด้วยอัลกอริทึม Support Vector Machine (SVM)

ตารางที่ 5 ผลการประเมินการเรียนรู้ด้วยอัลกอริทึม Support Vector Machine (SVM)

	Accuracy	Precision	Recall	F-Measure
2 Fold	80.69 %	81.19 %	80.69 %	80.94 %
3 Fold	81.27 %	81.53 %	81.16 %	81.34 %
4 Fold	82.91 %	83.28 %	82.79 %	83.03 %
5 Fold	82.45 %	82.95 %	82.31 %	82.63 %
6 Fold	83.91 %	84.23 %	83.80 %	84.01 %
7 Fold	83.91 %	84.46 %	83.76 %	84.11 %
8 Fold	83.55 %	83.94 %	83.42 %	83.68 %
9 Fold	83.18 %	83.93 %	83.01 %	83.46 %

	Accuracy	Precision	Recall	F-Measure
10 Fold	83.27 %	83.83 %	83.13 %	83.47 %

จากการทดลองการเรียนรู้ด้วยอัลกอริทึม Support Vector Machine (SVM) เพื่อจำแนกข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์โดยการทดสอบด้วยวิธีดังนี้

วิธี 2-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 80.69 ค่า Precision มีค่าอยู่ที่ร้อยละ 81.19 ค่า Recall มีค่าอยู่ที่ร้อยละ 80.69 และค่า F-measure มีค่าอยู่ที่ร้อยละ 80.94

วิธี 3-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 81.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 81.53 ค่า Recall มีค่าอยู่ที่ร้อยละ 81.16 และค่า F-measure มีค่าอยู่ที่ร้อยละ 81.34

วิธี 4-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 82.91 ค่า Precision มีค่าอยู่ที่ร้อยละ 83.28 ค่า Recall มีค่าอยู่ที่ร้อยละ 82.79 และค่า F-measure มีค่าอยู่ที่ร้อยละ 83.03

วิธี 5-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 82.45 ค่า Precision มีค่าอยู่ที่ร้อยละ 82.95 ค่า Recall มีค่าอยู่ที่ร้อยละ 82.31 และค่า F-measure มีค่าอยู่ที่ร้อยละ 82.63

วิธี 6-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 83.91 ค่า Precision มีค่าอยู่ที่ร้อยละ 84.23 ค่า Recall มีค่าอยู่ที่ร้อยละ 83.80 และค่า F-measure มีค่าอยู่ที่ร้อยละ 84.01

วิธี 7-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 83.91 ค่า Precision มีค่าอยู่ที่ร้อยละ 84.46 ค่า Recall มีค่าอยู่ที่ร้อยละ 83.76 และค่า F-measure มีค่าอยู่ที่ร้อยละ 84.11

วิธี 8-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 83.55 ค่า Precision มีค่าอยู่ที่ร้อยละ 83.94 ค่า Recall มีค่าอยู่ที่ร้อยละ 83.42 และค่า F-measure มีค่าอยู่ที่ร้อยละ 83.68

วิธี 9-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 83.18 ค่า Precision มีค่าอยู่ที่ร้อยละ 83.93 ค่า Recall มีค่าอยู่ที่ร้อยละ 83.01 และค่า F-measure มีค่าอยู่ที่ร้อยละ 83.46

วิธี 10-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ร้อยละ 83.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 83.83 ค่า Recall มีค่าอยู่ที่ร้อยละ 83.13 และค่า F-measure มีค่าอยู่ที่ร้อยละ 83.47

3.3 ผลการประเมินประสิทธิภาพการเรียนรู้ด้วยอัลกอริทึม K-Nearest Neighbour

ตารางที่ 6 ผลการประเมินประสิทธิภาพการเรียนรู้ด้วยอัลกอริทึม K-Nearest Neighbour

	Accuracy	Precision	Recall	F-Measure
2 Fold	74.18 %	76.18 %	73.87%	75.00 %
3 Fold	72.91 %	74.95 %	72.58 %	73.74 %
4 Fold	73.82 %	76.99 %	73.43 %	75.16 %
5 Fold	73.18 %	76.69 %	72.77 %	74.68 %
6 Fold	72.82 %	76.96 %	72.38 %	74.59 %
7 Fold	72.27 %	76.21 %	71.83 %	73.96 %
8 Fold	72.45 %	76.61 %	72.01 %	74.23 %
9 Fold	71.91 %	76.04 %	71.46 %	73.68 %
10 Fold	72.27 %	77.18 %	71.79 %	74.39 %

จากการทดลองการเรียนรู้ด้วยอัลกอริทึม K-Nearest Neighbour เพื่อจำแนกข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์โดยการทดสอบด้วยวิธีดังนี้

วิธี 2-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.18 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.18 ค่า Recall มีค่าอยู่ที่ร้อยละ 73.87 และค่า F-measure มีค่าอยู่ที่ร้อยละ 75.00

ที่ร้อยละ 73.87 และค่า F-measure มีค่าอยู่ที่ร้อยละ 75.00

วิธี 3-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 72.91 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.95 ค่า Recall มีค่าอยู่ที่ร้อยละ 72.58 และค่า F-measure มีค่าอยู่ที่ร้อยละ 73.74

วิธี 4-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 73.82 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.99 ค่า Recall มีค่าอยู่ที่ร้อยละ 73.43 และค่า F-measure มีค่าอยู่ที่ร้อยละ 75.16

วิธี 5-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 73.18 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.69 ค่า Recall มีค่าอยู่ที่ร้อยละ 72.77 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.68

วิธี 6-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 72.82 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.96 ค่า Recall มีค่าอยู่ที่ร้อยละ 72.38 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.59

วิธี 7-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 72.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.21 ค่า Recall มีค่าอยู่ที่ร้อยละ 71.83 และค่า F-measure มีค่าอยู่ที่ร้อยละ 73.96

วิธี 8-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 72.45 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.61 ค่า Recall มีค่าอยู่ที่ร้อยละ 72.01 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.23

วิธี 9-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 71.91 ค่า Precision มีค่าอยู่ที่ร้อยละ 76.04 ค่า Recall มีค่าอยู่ที่ร้อยละ 71.46 และค่า F-measure มีค่าอยู่ที่ร้อยละ 73.68

วิธี 10-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ร้อยละ 72.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 77.18 ค่า Recall มีค่าอยู่ที่ร้อยละ 71.79 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.39

ที่ร้อยละ 71.79 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.39

3.4 ผลการประเมินประสิทธิภาพการเรียนรู้ด้วยอัลกอริทึม Naive Bayes

ตารางที่ 7 ผลการประเมินประสิทธิภาพการเรียนรู้ด้วยอัลกอริทึม Naive Bayes

	Accuracy	Precision	Recall	F-Measure
2 Fold	73.64 %	73.71 %	73.55 %	73.63 %
3 Fold	74.36 %	74.39 %	74.30 %	74.34 %
4 Fold	75.82 %	75.88 %	75.74 %	75.81 %
5 Fold	74.27 %	74.29 %	74.22 %	74.25 %
6 Fold	74.64 %	74.71 %	74.55 %	74.63 %
7 Fold	74.09 %	74.12 %	74.03 %	74.07 %
8 Fold	74.55 %	74.60 %	74.47 %	74.53 %
9 Fold	74.27 %	74.35 %	74.19 %	74.27 %
10 Fold	75.27 %	75.32 %	75.20 %	75.26 %

จากการทดลองการเรียนรู้ด้วยอัลกอริทึม Naive Bayes เพื่อจำแนกข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์โดยทำการทดสอบด้วยวิธีดังนี้

วิธี 2-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 73.64 ค่า Precision มีค่าอยู่ที่ร้อยละ 73.71 ค่า Recall มีค่าอยู่ที่ร้อยละ 73.55 และค่า F-measure มีค่าอยู่ที่ร้อยละ 73.63

วิธี 3-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.36 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.39 ค่า Recall มีค่าอยู่ที่ร้อยละ 74.30 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.34

วิธี 4-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 75.82 ค่า Precision มีค่าอยู่ที่ร้อยละ 75.88 ค่า Recall มีค่าอยู่ที่ร้อยละ 75.74 และค่า F-measure มีค่าอยู่ที่ร้อยละ 75.81

ที่ร้อยละ 75.74 และค่า F-measure มีค่าอยู่ที่ร้อยละ 75.81

วิธี 5-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.29 ค่า Recall มีค่าอยู่ที่ร้อยละ 74.22 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.25

วิธี 6-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.64 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.71 ค่า Recall มีค่าอยู่ที่ร้อยละ 74.55 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.63

วิธี 7-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.09 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.12 ค่า Recall มีค่าอยู่ที่ร้อยละ 74.03 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.07

วิธี 8-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.55 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.60 ค่า Recall มีค่าอยู่ที่ร้อยละ 74.47 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.53

วิธี 9-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ ร้อยละ 74.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 74.35 ค่า Recall มีค่าอยู่ที่ร้อยละ 74.19 และค่า F-measure มีค่าอยู่ที่ร้อยละ 74.27

วิธี 10-fold cross validation สามารถประเมินประสิทธิภาพได้ดังนี้ ค่าAccuracy มีค่าอยู่ที่ร้อยละ 75.27 ค่า Precision มีค่าอยู่ที่ร้อยละ 75.32 ค่า Recall มีค่าอยู่ที่ร้อยละ 75.20 และค่า F-measure มีค่าอยู่ที่ร้อยละ 75.26

3.5 ผลการเปรียบเทียบประเมินประสิทธิภาพระหว่าง อัลกอริธึม Support Vector Machine (SVM) , K-Nearest Neighbour และ อัลกอริธึม Naive Bayes

จากผลการทดลองในการเรียนรู้โดยการทดสอบด้วยวิธี 6-fold cross validation อัลกอริธึม Support Vector Machine (SVM) ให้ค่า Accuracy , Precision , Recall และ F-Measure ที่ร้อยละ 83.91, 84.23, 83.80 และ 84.01 ตามลำดับ ซึ่งสูงอัลกอริธึม Naive Bayes ที่ให้ค่า

Accuracy, Precision, Recall และ F-Measure ที่ร้อยละ 74.64, 74.71, 74.55 และ 74.63 ตามลำดับ และยังสูงกว่าอัลกอริธึม K-Nearest Neighbour ที่ให้ค่า Accuracy, Precision, Recall และ F-Measure ที่ร้อยละ 72.82, 76.96, 72.38 และ 74.59 ตามลำดับ ดังตารางที่ 8

ตารางที่ 8 ผลการเปรียบเทียบประเมินประสิทธิภาพระหว่างอัลกอริธึม Support Vector Machine (SVM) , K-Nearest Neighbour และ อัลกอริธึม Naive Bayes

	Support Vector Machine	Naive Bayes	K-Nearest Neighbour
Accuracy	83.91 %	75.82 %	74.18 %
Precision	84.23 %	75.88 %	76.18 %
Recall	83.80 %	75.74 %	73.87 %
F-Measure	84.01 %	75.81 %	75.00 %

4. อภิปรายผลและสรุป

4.1 สรุปผลการวิจัย

งานวิจัยนี้ได้นำเสนอการพัฒนาแบบจำลองในการตรวจจับข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์ โดยใช้วิธีซัพพอร์ตเวกเตอร์แมชชีน ในการพัฒนาแบบจำลองเริ่มต้นจากรวบรวมข้อมูล วิเคราะห์ข้อมูลที่รวบรวม การตัดคำภาษาไทย และการกำจัดคำหยุด จากนั้นพัฒนาแบบจำลองเพื่อจำแนกข้อมูล และเปรียบเทียบอัลกอริธึมเพื่อประเมินประสิทธิภาพในการจำแนกข้อความทั้งในด้านความถูกต้อง (Accuracy) ความแม่นยำ (Precision) ความระลึก (Recall) และค่าความถ่วงดุล (F-Measure) ซึ่งประกอบด้วยอัลกอริธึม Support Vector Machine (SVM) อัลกอริธึม K-Nearest Neighbour และ อัลกอริธึม Naive Bayes ซึ่งพบว่า อัลกอริธึม Support Vector Machine (SVM) มีประสิทธิภาพในการจำแนกข้อความมากที่สุด ซึ่งมีค่าAccuracy อยู่ที่ร้อยละ 83.91 ค่า Precision อยู่ที่ร้อยละ 84.23 ค่า Recall อยู่ที่ร้อยละ 83.80 และค่า F-measure อยู่ที่ร้อยละ 84.01 ซึ่งเป็นไปตามสมมติฐานการวิจัยที่ตั้งไว้

4.2 อภิปรายผลการวิจัย

อัลกอริทึม Support Vector Machine (SVM) มีประสิทธิภาพในการตรวจจับข้อความภาษาไทยที่เป็นการกลั่นแกล้งทางไซเบอร์มากที่สุดซึ่งสอดคล้องกับกับงานวิจัยของ Noviantho, Isa และ Ashianti [5] ที่ชี้ให้เห็นว่าอัลกอริทึม Support Vector Machine (SVM) มีประสิทธิภาพในการจำแนกมากกว่าอัลกอริทึม Naive Bayes และงานวิจัยของของ Nurrahmi และ Nurjanah [24] ที่ชี้ให้เห็นว่าอัลกอริทึม Support Vector Machine (SVM) มีประสิทธิภาพในการจำแนกมากกว่าอัลกอริทึม K-Nearest Neighbour ทั้งนี้เนื่องจากอัลกอริทึม Support Vector Machine (SVM) ใช้สมการเชิงเส้นที่ทำการแบ่งข้อมูลออกเป็นสองฝั่ง จึงสามารถที่จะจำแนกกลุ่มของข้อมูลที่มีจำนวนเพียง 2 กลุ่มได้ดีกว่าอัลกอริทึม K-Nearest Neighbour ที่ใช้การวัดระยะห่างของข้อมูล และอัลกอริทึม Naive Bayes ที่ใช้ความน่าจะเป็น โดยที่ผลการประเมินประสิทธิภาพของอัลกอริทึม Support Vector Machine (SVM) มีค่าความถูกต้องอยู่ที่ร้อยละ 83.91 ซึ่งอยู่ในเกณฑ์ดีตรงตามสมมติฐานแต่ยังถือว่ามีความถูกต้องที่น้อยอยู่เนื่องจากจำนวนข้อมูลที่ใช้ในงานวิจัยครั้งนี้ยังมีจำนวนที่น้อยอยู่ รวมถึงการใช้คำภาษาไทยในปัจจุบันนั้นมีโครงสร้างไม่แน่นอน มีความหลากหลาย และความยืดหยุ่นของการใช้ภาษาไทย ซึ่งสอดคล้องกับงานวิจัยของ กานดาแผ้วพานกุล [36] ที่ให้ข้อแนะนำเดียวกันนี้ จึงทำให้เกิดความผิดพลาดในการจำแนกได้

4.3 ข้อจำกัดในงานวิจัย

4.3.1 สื่อสังคมออนไลน์บางตัวไม่อนุญาตให้ใช้ API ข้อความแสดงความคิดเห็นได้ จึงจำเป็นต้องใช้การดึงข้อความแสดงความคิดเห็นด้วยตัวเอง ซึ่งทำให้เกิดความล่าช้าในการรวบรวมข้อมูล

4.3.2 การตัดคำภาษาไทยยังมีความผิดพลาด เพราะข้อความจากสื่อออนไลน์มีการใช้คำที่ไม่ตรงหลักภาษาไทย ซึ่งทำให้ไม่สามารถวิเคราะห์ข้อความนั้นได้อย่างถูกต้อง และเกิดความผิดพลาดในการวัดประสิทธิภาพของอัลกอริทึม

4.4 ข้อเสนอแนะในการต่อยอดงานวิจัย

การต่อยอดในอนาคตอาจสามารถนำไปพัฒนาต่อยอดในการป้องกันการกลั่นแกล้งทางไซเบอร์บนสื่อโซเชียลออนไลน์

ไลน์ได้ ซึ่งจะช่วยให้สังคมออนไลน์นั้นปลอดภัยและน่าอยู่มากยิ่งขึ้น

5. เอกสารอ้างอิง

- [1] เด็กไทยเสี่ยงภัยออนไลน์ ห่วง Cyber bullying สูงกว่าค่าเฉลี่ยโลก. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 1 ตุลาคม 2561]. เข้าถึงได้จาก: <https://www.it24hrs.com/2018/cyber-bullying-thai-child-above-the-world-average/>
- [2] คุณรู้จักความหมายของ Cyberbullying ดีพอแล้วหรือยัง?. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 1 ตุลาคม 2561]. เข้าถึงได้จาก: <https://medium.com/@supanatwongchai/คุณรู้จักความหมายของ-cyberbullying-ดีพอแล้วหรือยัง-ce4d6f434b47>
- [3] Stop Bullying หยุดกลั่นแกล้งทางโลกออนไลน์ หลังไทยติดอันดับ Top5 ของโลก!. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 1 ตุลาคม 2561]. เข้าถึงได้จาก: <https://brandinside.asia/stop-bullying-thailand-top5/>
- [4] ทำความรู้จัก Stop Cyberbullying Day วันหยุดการกลั่นแกล้งบนโลกออนไลน์. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 1 ตุลาคม 2561]. เข้าถึงได้จาก: <http://stopbullying.lovecarestation.com/articles-news/articles/stop-cyberbullying-day/>
- [4] Noviantho, Sani Muhamad Isa, Livia Ashianti. Cyberbullying classification using text mining. 2017 1st International Conference on Informatics and Computational Sciences (ICICoS); 2017. 241-246.
- [6] รวิสุตา เทศเมือง นิเวศ จิระวิจิตชัย. การวิเคราะห์ความคิดเห็นภาษาไทยเกี่ยวกับการรีวิวสินค้าออนไลน์ โดยใช้ขั้นตอนวิธีฟัฟฟอร์ดเวกเตอร์แมทซิน. Engineering Journal of Siam University. ปีที่ 18 ฉบับที่.34; 2017
- [7] cybe. [ออนไลน์]; 2560. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <https://dict.longdo.com/search/-cyber->
- [8] ป้องภัยใกล้ตัวลูกจาก Cyber bullying. [ออนไลน์]; 2560. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก:

- <https://www.dmh.go.th/news-dmh/view.asp?id=27717>
- [9] ทำความรู้จักเกี่ยวกับ Social Media และผู้ให้บริการ Social Media bullying. [ออนไลน์]; 2557. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <https://medium.com/@Natthapete/ทำความรู้จักเกี่ยวกับ-social-media-และผู้ให้บริการ-social-media-f210a83d2d76>
- [10] การทำเหมืองข้อมูล (Data Mining). [ออนไลน์]; 2560. [เข้าถึงเมื่อ 1 ตุลาคม 2561]. เข้าถึงได้จาก: <http://compcenter.bu.ac.th/news-information/data-mining>
- [11] อัลกอริทึม คืออะไร มีความสำคัญอย่างไร การเขียนโปรแกรม. [ออนไลน์]; 2562. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <https://www.mindphp.com/คู่มือ/73-คืออะไร/6842-what-is-a-algorithm.html>
- [12] API คืออะไร เอพีไอ คือ ช่องทางหนึ่งที่จะเชื่อมต่อกับเว็บไซต์ผู้ให้บริการ API จากที่อื่น. [ออนไลน์]; 2560. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <https://www.mindphp.com/คู่มือ/73-คืออะไร/2038-api-คืออะไร.html>
- [13] เหมืองข้อความและการประยุกต์ใช้. [ออนไลน์]; 2556. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <http://www.tci-thaijo.org/index.php/pimjournal/article/view/12048>
- [14] อัลกอริทึม Support Vector Machine (SVM). [ออนไลน์]; 2560. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <http://kokzard.blogspot.com/2011/10/jfjkdshfkjsldf.html>
- [15] ขั้นตอนวิธีการค้นหาเพื่อนบ้านใกล้สุด k ตัว. [ออนไลน์]; 2560. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก https://th.wikipedia.org/wiki/ขั้นตอนวิธีการค้นหาเพื่อนบ้านใกล้สุด_k_ตัว.
- [16] การเรียนรู้แบบเบย์ (Bayesian Learning). [ออนไลน์]; 2558. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <https://mis.csit.sci.tsu.ac.th/noppamas/download/DataMining>.
- [17] Text normalization. [ออนไลน์]; 2560. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: https://en.wikipedia.org/wiki/Text_normalization
- [18] Virach Sornlertlamvanich. Word segmentation for Thai in machine translation system. [ออนไลน์]; 1993. เข้าถึงได้จาก https://www.researchgate.net/profile/VirachSornlertlamvanich/publication/243659316_Word_segmentation_for_Thai_in_machine_translation_system/links/5451ff2e0cf2bf864cbabb07/Word-segmentation-for-Thai-in-machine-translation-system.pdf
- [19] นิเวศ จิระวิชิตชัย ปริญา สงวนสัตย์ และพยุง มีสัจ. การพัฒนาประสิทธิภาพการจัดหมวดหมู่เอกสารภาษาไทยแบบอัตโนมัติ. NIDA Development Journal. Vol.51 No.3; 2011.
- [20] Salton, G. and C. Buckley. Term-weighting approaches in automatic text retrieval. Information Processing and Management. 24(5); 1988. 513-523.
- [21] การแบ่งข้อมูลเพื่อนำทดสอบประสิทธิภาพของโมเดล. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <http://dataminingtrend.com/2014/data-mining-techniques/cross-validation/>
- [22] Cyberbullying ภัยใกล้ตัวบนโลกออนไลน์ที่ไม่ควรมองข้าม. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 4 มีนาคม 2562]. เข้าถึงได้จาก: <https://theprotype.pim.ac.th/2018/01/04/whatiscyberbullying/>
- [23] Selma Ayúe Özel, Esra Saraç, Seyran Akdemir Hülya Aksu. Detection of Cyberbullying on Social Media Messages in Turkish. 2017 International Conference on Computer Science and Engineering (UBMK); 2017.
- [24] Hani Nurrahmi, Dade Nurjanah. Indonesian Twitter Cyberbullying Detection using Text Classification and User Credibility. 2018 International Conference on Information and Communications Technology (ICOIAC); 2018.

- [25] Kelly Reynolds, April Kontostathis, Lynne Edwards. (Using Machine Learning to Detect Cyberbullying. 2011 10th International Conference on Machine Learning and Applications and Workshops; 2011.
- [26] Batoul Haidar, Maroun Chamoun, Ahmed Serhrouchni. Multilingual Cyberbullying Detection System Detecting Cyberbullying in Arabic Content. 2017 1st Cyber Security in Networking Conference (CSNet); 2017.
- [27] พิมพ์วัลย์ บุญมงคล ญัฐรัชต์ สาเมาะ มุจลินท์ ชลรัตน์. การพัฒนาศักยภาพเครือข่ายโรงเรียนเพื่อลดความรุนแรงในพื้นที่. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 10 ตุลาคม 2561]. เข้าถึงได้จาก: <http://www.sh.mahidol.ac.th/chps/th/3175>.
- [28] Silvana Tambosi, Vanessa Mondini, Gustavo Borges, Maria José Domingues. 2015. Cyberbullying: Concerns of Teachers and School Involvement. 2015 Ninth International Conference on Complex, Intelligent, and Software Intensive Systems.
- [29] วัฒนวรรณ อินทรพล วุฒิชัย ร่มสายหยุด. [ออนไลน์] 2561. [สืบค้นวันที่ 10 ตุลาคม 2561]. การทำเหมืองข้อความสำหรับการแก้ปัญหาข้อร้องเรียนเกี่ยวกับการบริการลูกค้าที่บริษัทการสื่อสาร. <https://tci-thaijo.org/index.php/itm-journal/article/download/115352/89152/>.
- [30] อมรทิพย์ อมราภิบาล. สายหยุด. เหยื่อการรังแกผ่านโลกไซเบอร์ในกลุ่มเยาวชน: ปัจจัยเสี่ยงผลกระทบต่อสุขภาพจิตและการปรึกษาบุคคลที่สาม. [ออนไลน์]; 2561. [เข้าถึงเมื่อ 10 ตุลาคม 2561]. เข้าถึงได้จาก: <https://www.tci-thaijo.org/index.php/RMCS/article/view/59642>.
- [31] พิษณุสินี กิจวัฒนาถาวร ธรา อังสกุล จิตมินต์ อังสกุล. ระบบสกัดความรู้จากบทวิจารณ์โรงแรมออนไลน์โดยใช้ตรรกศาสตร์คลุมเครือ. [ออนไลน์]; 2556. [เข้าถึงเมื่อ 10 ตุลาคม 2561]. เข้าถึงได้จาก: <https://www.tci-thaijo.org/index.php/kmutnb-journal/article/view/18234>.
- [32] สุริยศักดิ์ เลิศสกุลสมบูรณ์. การสกัดความสัมพันธ์แบบสัณยภาพจากเอกสารงานวิจัยทางวิทยาศาสตร์. [ออนไลน์]; 2556. [เข้าถึงเมื่อ 10 ตุลาคม 2561]. เข้าถึงได้จาก: <http://libdoc.dpu.ac.th/thesis/149444.pdf>.
- [33] ศิริพร อ่วมมีเพียร สันติพงษ์ ไทยประยูร. ระบบติดตามการคัดลอกเนื้อหาเว็บไซต์อัตโนมัติ โดยใช้วิธีการเลือกข้อความสำคัญ. วารสารเทคโนโลยีสารสนเทศ. ปีที่ 12 ฉบับที่ 2; 2559
- [34] Yasin N. Silva, Christopher Rich, Deborah Hall. BullyBlocker: Towards the Identification of Cyberbullying in Social Networking Sites. 2016 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM); 2016. 1377-1379.
- [35] I-Hsien Ting Wun Sheng Liou Dario Liberona Shyue-Liang Wang Giovanni Mauricio Tarazona Bermudez. Towards the Detection of Cyberbullying Based on Social Network Mining Techniques. 2017 International Conference on Behavioral, Economic, Socio-cultural Computing (BESCom); 2017.
- [36] กานดา แผ้ววัฒนากุล. การวิเคราะห์เหมืองข้อเสนอแนะจากบทวิจารณ์รายการโทรทัศน์. [ออนไลน์]; 2555. [เข้าถึงเมื่อ 10 ตุลาคม 2561]. เข้าถึงได้จาก: <http://libdocms.nida.ac.th/thesis6/2555/b179815.pdf>.