

Research Article

เทคนิคทางปัญญาประดิษฐ์สำหรับการตรวจข่าวปลอมภาษาไทย

Artificial Intelligent Techniques for Thai Fake News Detection

พยุ่ง มีสัจ^{1*}, พนม คลีฉายา², อองอาจ อุ่นอนันต์³, กมลรัตน์ กิจรุ่งไพศาล⁴

Phayung Meesad^{1*}, Phnom Kleechaya², Aongart Aun-a-nan³, Kamonrat Kijrungsaisarn⁴

^{1,3}คณะเทคโนโลยีสารสนเทศและนวัตกรรมดิจิทัล มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ ประเทศไทย

^{1,3}Faculty of Information Technology and Digital Innovation, King Mongkut's University of Technology North Bangkok, Thailand

^{2,4}หน่วยปฏิบัติการวิจัยและพัฒนาทรัพยากรมนุษย์ด้านความรู้ทางดิจิทัลและการรู้เท่าทันสื่อ (DIRU) จุฬาลงกรณ์มหาวิทยาลัย
ประเทศไทย

^{2,4}Digital Intelligence and Literacy Research Unit (DIRU), Chulalongkorn University, Thailand

*E-mail: phayung.m@itd.kmutnb.ac.th

Received: 13/06/2021; Revised: 24/10/2021; Accepted: 23/04/2022

บทคัดย่อ

ด้วยเทคโนโลยีดิจิทัลที่ทันสมัยเพิ่มความสะดวกแก่ผู้ใช้งานในการเผยแพร่ข่าวสารบนสื่อออนไลน์ โดยส่วนหนึ่งได้เผยแพร่ข่าวปลอมที่มีจำนวนมากซึ่งเป็นปัญหาทำให้ผู้หลงเชื่อข่าวปลอมว่าเป็นข่าวจริงและนำไปแชร์ต่อ ทำให้เกิดปัญหาอื่นๆ ต่อเนื่องตามมา ปัจจุบันยังไม่มีเครื่องมือที่จะช่วยตรวจสอบข่าวปลอมภาษาไทย ที่จะช่วยชะลอการแพร่กระจายข่าว การตรวจจับข่าวปลอมแบบอัตโนมัติถือว่าเป็นงานที่ยากและมีความท้าทายสูง เนื่องจากข่าวมีความเคลื่อนไหวอย่างมาก งานวิจัยนี้เสนอวิธีการตรวจสอบข่าวปลอมภาษาไทยด้วยเทคนิคปัญญาประดิษฐ์ ผู้วิจัยใช้สามเทคนิคหลักในการสร้างแบบจำลองการตรวจจับข่าวปลอม ได้แก่ การสืบค้นข้อมูลสารสนเทศ การประมวลผลภาษาธรรมชาติ และการเรียนรู้ของเครื่อง การสืบค้นข้อมูลสารสนเทศเป็นกลไกในการดึงข้อมูลจากเว็บไซต์ข่าวออนไลน์และโซเชียลมีเดีย การประมวลผลภาษาธรรมชาติจะวิเคราะห์ข่าวที่ได้กลับคืนมาเพื่อสกัดข้อมูลคุณลักษณะมีความโดดเด่นเป็นอย่างดี การเรียนรู้ของเครื่องทำการรับข้อมูลฟีดแบ็กและจัดประเภทบทความข่าวออกเป็นสามประเภท ได้แก่ ข่าวจริง ข่าวปลอม และข่าวน่าสงสัย ในการเตรียมข้อมูลสอน ผู้วิจัยใช้เว็บครอว์เลอร์เพื่อรวบรวมจัดเก็บข้อมูลจำนวน 53,220 ตัวอย่าง โดยทำการจัดประเภทไว้ล่วงหน้าเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย เพื่อป้องกันความโน้มเอียง จำนวนตัวอย่างข้อมูลในแต่ละกลุ่ม สุ่มเลือกให้มีความสมดุล ผู้วิจัยใช้วิธีการสอนแบบจำลองแบบไขว้ทดสอบข้อมูล 10 ชุด แบบจำลองการเรียนรู้ของเครื่องที่ใช้ในการศึกษา ได้แก่ การถดถอยโลจิสติกส์ เพื่อนบ้านใกล้เคียง นาอีฟเบย์ โครงข่ายประสาทเทียมชนิดเพอร์เซ็ปตรอนหลายชั้น ซัพพอร์ต

เวกเตอร์แมชชีน ป่าสุ่ม และโครงข่ายประสาทเทียมชนิดหน่วยความจำระยะสั้นแบบยาว ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยด้านค่าความถูกต้อง ค่าความแม่นยำ ค่าการเรียกคืน และค่าถ่วงดุล ของการเรียนรู้ของเครื่องชนิดต่าง ๆ กลุ่มโมเดลที่พบว่าดีที่สุดในงานวิจัยนี้ได้แก่ DT, RF, SVM, LSTM และ MLP ซึ่งเหมาะสำหรับการเรียนรู้ของเครื่องที่จะนำไปใช้เป็นระบบตรวจข่าวปลอม

คำสำคัญ: ข่าวปลอม, ปัญญาประดิษฐ์, การสืบค้นสารสนเทศ, การประมวลผลภาษาธรรมชาติ, การเรียนรู้ของเครื่อง

Abstract

With modern digital technology, it is more convenient for users to distribute news on online media. Some of them have spread a lot of fake news, which is a problem causing people to believe fake news as real news and share it with others. Currently, there is no tool to detect fake news and interrupt the spread of Thai fake news. Detecting fake news automatically is a difficult and challenging task due to the dynamics of the news. This research presents a method for detecting fake news in Thai with artificial intelligence techniques based on information retrieval, natural language processing, and machine learning. In preparing the training data, the researchers used web crawlers to collect 53,220 samples and pre-categorize them as real, fake, and suspicious news. We balanced the number of data samples in each group to avoid biasing. We trained machine learning models based on 10-fold cross validation. The machine learning models used in the study were Logistic Regression (LR), K-Nearest Neighbor (KNN), Naïve Bayesian (NB), Multilayer Perceptron (MLP), Support Vector Machine (SVM), Random Forest (RF) and Short-Term Long-Term Memory (LSTM). We compared the average performance of different types of machine learning based on accuracy, precision, recall, and f-measure. The models found best in this study were DT, RF, SVM, LSTM, and MLP, suitable for machine learning a fake news detection system.

Keywords: fake news, artificial intelligence, information retrieval, natural language processing, machine learning

บทนำ

ด้วยเทคโนโลยีสารสนเทศหรือดิจิทัลในปัจจุบัน การสื่อสารข้อมูลสามารถส่งถึงกันได้อย่างรวดเร็ว วัฒนาการของเทคโนโลยีสารสนเทศและการสื่อสารได้เพิ่มจำนวนผู้ใช้อินเทอร์เน็ตอย่างมากซึ่งเปลี่ยนวิธีการใช้ข้อมูลและข่าวสารจากแบบดั้งเดิมเป็นแบบดิจิทัล ทำให้เกิดความสะดวกรวดเร็วทั้งผู้เสนอข่าวและผู้อ่านข่าว ในความสะดวกรวดเร็วของระบบอินเทอร์เน็ตทำให้เกิดข่าวปลอมขึ้นจำนวนมากเช่นกัน ความนิยมของสาธารณชนในโซเชียลมีเดีย และเครือข่ายสังคม ทำให้เกิดการแพร่กระจายของข่าวปลอม โดยที่การบิดเบือน และมุมมองที่รุนแรง

ข่าวปลอมได้กลายเป็นหนึ่งในความกังวลที่สำคัญเนื่องจากมีศักยภาพที่จะทำให้รัฐบาลไม่มั่นคง หรืออาจทำให้เป็นอันตรายต่อสังคม ด้วยสื่อสมัยใหม่ ข่าวปลอมอาจกล่าวได้ว่าเป็นภัยคุกคามที่สำคัญ ส่งผลให้ความเชื่อมั่นของรัฐบาลลดลง กระทบประชาชนในทุก ๆ ด้าน

การตรวจจับและลดผลกระทบของข่าวปลอมเป็นหนึ่งในปัญหาพื้นฐานของยุคสมัยปัจจุบันและได้รับความสนใจอย่างกว้างขวาง ข่าวปลอมเป็นหัวข้อเกี่ยวข้องกับหลายสาขาซึ่งนักวิชาการสาขาต่าง ๆ ได้ทำการศึกษาแง่มุมต่าง ๆ ของข่าวปลอม เช่น การเรียนรู้ของเครื่อง ฐานข้อมูล วารสารศาสตร์ รัฐศาสตร์ และอื่น ๆ อีกมากมาย (Lakshmanan *et al.*, 2019) ข่าวปลอมที่เผยแพร่อยู่บนอินเทอร์เน็ตมีข้อมูลในรูปแบบที่หลากหลาย เช่น เอกสาร วิดีโอ และไฟล์เสียง ข่าวปลอมออนไลน์ในรูปแบบที่ไม่มีโครงสร้าง จึงค่อนข้างยากที่จะตรวจจับและจำแนกเนื่องจากต้องใช้ความเชี่ยวชาญของมนุษย์อย่างเคร่งครัด อย่างไรก็ตาม เทคนิคการคำนวณเชิงปัญญาประดิษฐ์ เช่น การประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) และการเรียนรู้ของเครื่อง (Machine Learning: ML) สามารถนำมาประยุกต์ใช้ตรวจจับและจำแนกข้อความหรือบทความที่เป็นประเภทหลอกลวงโดยธรรมชาติจากบทความที่อิงข้อเท็จจริง (Ahmad *et al.*, 2020) ในช่วงหลายปีที่ผ่านมาได้มีนักวิจัยเกี่ยวกับปัญญาประดิษฐ์ที่นำเสนอการตรวจข่าวปลอมจำนวนมาก

Granik & Mesyura (2017) ใช้วิธีการง่าย ๆ สำหรับการตรวจข่าวปลอมโดยใช้ตัวแยกประเภทนาอิวเบย์ (Naïve Baye) วิธีนี้ใช้เป็นระบบซอฟต์แวร์และทดสอบกับชุดข้อมูลของโพสต์ข่าวบนเฟซบุ๊ก โดยมีความแม่นยำในการจำแนกประเภทที่ประมาณ 74% ในชุดทดสอบ ซึ่งเป็นผลลัพธ์ที่ดีเมื่อพิจารณาจากความเรียบง่ายสัมพัทธ์ของแบบจำลอง อีกทั้ง Shu *et al.* (2019) ได้ศึกษาและนำเสนอวิธีการต่าง ๆ เพื่อจะตรวจสอบคุณสมบัติข่าวปลอมและวิธีการแพร่กระจายข่าวทางเครือข่ายสังคมออนไลน์ แนะนำประเภทเครือข่ายยอดนิยม และเสนอว่าเครือข่ายเหล่านี้สามารถใช้ในการตรวจจับและบรรเทาข่าวปลอมในสื่อสังคมออนไลน์ได้

สำหรับเทคนิคใหม่ ๆ Shishah (2021) ใช้เทคนิคตัวแทนเข้ารหัสแบบสองทิศทางจากตัวแปลง (Bidirectional Encoder Representations from Transformers: BERT) การเรียนรู้ร่วมกันที่รวมการจำแนกคุณสมบัติเชิงสัมพันธ์ (Relational Features Classification: RFC) และความสัมพันธ์ชื่อนิบุคคล (Name Entity Relation: NER) การทดลองกับชุดข้อมูลจริงสองชุด พบว่าประสิทธิภาพเป็นที่น่าพอใจ โดยวัดด้วยความถูกต้อง (Accuracy) ค่าถ่วงดุล (F-Measure) และพื้นที่ใต้เส้นโค้ง (Area Under Curve: AUC) เอกลักษณะของกรอบการทำงานร่วมกันทำให้น้ำหนักที่มีความหมายต่อคุณลักษณะ ซึ่งนำไปสู่ประสิทธิภาพที่ดีขึ้นเมื่อเทียบกับเทคนิคปัญญาประดิษฐ์อื่น ๆ นอกจากนี้ Birunda & Devi (2021) ได้เสนอการตรวจข่าวปลอมจากแหล่งข่าวหลายแหล่งแบบอัตโนมัติ โดยแยกคุณลักษณะแบบข้อความจากบทความข่าวจริงและข่าวปลอมโดยใช้ค่าน้ำหนักซึ่งเป็นการถ่วงและค่าถ่วงเอกสารกลับด้าน (Term Frequency - Inverted Document Frequency: TF-IDF) และใช้คะแนนความน่าเชื่อถือของแหล่งที่มาคำนวณตามคุณลักษณะไชด์ยูอาร์แอล ร่วมกับคุณสมบัติแบบข้อความกับคะแนนความน่าเชื่อถือของแหล่งที่มาหลายแหล่ง ความน่าเชื่อถือของข่าวจะถูกประมาณการ กรอบการทำงานที่เสนอนำไปใช้กับตัวแยกประเภทการ

เรียนรู้ของเครื่องเพื่อตรวจสอบประสิทธิภาพในการตรวจหาข่าวปลอม ผลการทดลองได้ประสิทธิภาพของกรอบงานที่เสนอด้วยอัลกอริทึม Gradient Boosting ประมาณ 99.5% จนถึงระดับสูงสุด (Birunda & Devi, 2021) ในการตรวจจับข่าวปลอมภาษาอื่นๆ ที่ไม่ใช่ภาษาอังกฤษ DongHo Lee *et al.* (2019) เสนอสถาปัตยกรรมการเรียนรู้เชิงลึกสำหรับการตรวจจับข่าวปลอมที่เขียนเป็นภาษาเกาหลี โดยใช้สถาปัตยกรรมการเรียนรู้เชิงลึก และใช้รูปแบบการฝังคำที่เรียนรู้โดยหน่วยพยางค์ หลังจากการสอนและทดสอบการใช้งาน ระบบมีความถูกต้องที่มีความหมายสำหรับการจัดประเภทเนื้อหาและความคลาดเคลื่อนของบริบท แต่ความแม่นยำยังต่ำสำหรับการจำแนกประเภทพาดหัวและความคลาดเคลื่อนของเนื้อหา

ในประเทศไทยเองก็เริ่มมีงานวิจัยเกี่ยวกับการตรวจจับข่าวปลอมของไทยใช้เทคนิคปัญญาประดิษฐ์ Aphiwongsophon & Chongstitvatana (2018) ทำการศึกษาข่าวปลอมภาษาไทยจาก Twitter โดยใช้วิธีการยอคนิยมสามวิธีในการทดลอง ได้แก่ นาอ็อบเบย์ (Naive Bayes: NB) โครงข่ายประสาท (Neural Network: NN) และซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) การศึกษาพบว่ากระบวนการทำให้ข้อมูลเป็นมาตรฐานเป็นสิ่งจำเป็นสำหรับการทำความสะอาดข้อมูลก่อนที่จะป้อนข้อมูลไปสู่แบบจำลองการเรียนรู้ของเครื่องเพื่อจำแนกประเภทข้อมูล แบบจำลองที่ดีที่สุดคือ SVM มีความถูกต้องถึง 99.90% ซึ่งงานวิจัยนี้เน้นที่ข้อมูล Twitter เท่านั้น ไม่มีการใช้งานสำหรับการตรวจจับออนไลน์ นอกจากนี้ Mookdarsanit & Mookdarsanit (2021) เสนอกรอบการเรียนรู้เชิงลึกสำหรับการตรวจจับข่าวปลอมของไทยเกี่ยวกับโควิด-19 โดยใช้ข้อความจากโซเชียล งานวิจัยนี้ได้สร้างแบบจำลองการเรียนรู้แบบถ่ายโอน ซึ่งรวมถึง BERT, Universal Language Model Fine-Tuning (ULMFIT) และ Generative Pre-trained Transformer (GPT) นักวิจัยใช้ชุดข้อมูลเปิดข่าวโควิด-19 ที่แปลจากภาษาอังกฤษเป็นภาษาไทยและเตรียมการสอนแบบจำลองการเรียนรู้เชิงลึกเกี่ยวกับโควิด-19 เพื่อปรับแต่งชุดข้อมูลให้เข้ากับประเทศไทย โดยนักวิจัยใช้ข้อมูลเพิ่มเติมโดยรวบรวมข้อมูลข้อความภาษาไทยจากโซเชียลมีเดียและระบุว่าเป็นประเภทข่าวปลอมและข่าวจริง ผลลัพธ์ที่ดีที่สุดจากการทดลองทำให้ประสิทธิภาพความถูกต้องดีที่สุดถึง 72.93% ไม่มีรายงานกรณีการใช้งานจริงสำหรับข่าวปลอมของไทยในการวิจัยนี้

สำหรับงานวิจัยครั้งนี้ผู้วิจัยเสนอวิธีทางปัญญาประดิษฐ์สำหรับการตรวจข่าวปลอมภาษาไทย วิธีที่นำเสนอในงานวิจัยนี้ผู้วิจัยเสนอกรอบการทำงานเพื่อสร้างระบบตรวจจับข่าวปลอมออนไลน์ของไทยโดยอัตโนมัติ กรอบงานที่เสนอประกอบด้วยสามโมดูล ได้แก่ การสืบค้นข้อมูลสารสนเทศ การประมวลผลภาษาธรรมชาติ และการเรียนรู้ของเครื่อง งานวิจัยนี้มีวัตถุประสงค์หลัก คือ 1) นำเสนอกรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์ 2) นำเสนอลักษณะเด่นจากการวิเคราะห์ภาษาธรรมชาติ 3) นำเสนอแบบจำลองการเรียนรู้ของเครื่องสำหรับการตรวจจับข่าวปลอมภาษาไทย และ 4) พัฒนาเว็บแอปพลิเคชันตรวจข่าวปลอมภาษาไทย

วรรณกรรมที่เกี่ยวข้อง

หลายปีที่ผ่านมา มีนักวิจัยนำเสนอการตรวจข่าวปลอมด้วยเทคนิคด้านปัญญาประดิษฐ์มากมายหลายกลุ่มด้วยกัน โดยส่วนมากใช้เทคนิคที่เกี่ยวข้องกับปัญญาประดิษฐ์ส่วนมากอยู่ในระดับการประยุกต์ใช้การเรียนรู้ของเครื่อง ซึ่งเป็นการสอนข้อมูลให้เครื่องรู้รูปแบบของข่าว ในการที่จะสอนข้อมูลข้อความข่าวให้เครื่องรู้จำนั้นต้องใช้ศาสตร์ที่เกี่ยวข้องหลัก ๆ คือการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) (Ayoub et al., 2021) ซึ่งเป็นปัญญาประดิษฐ์ (Artificial Intelligence: AI) ที่เกี่ยวข้องกับการให้คอมพิวเตอร์สามารถเข้าใจข้อความและคำพูดในลักษณะเดียวกับที่มนุษย์สามารถทำได้ เป็นกระบวนการวิเคราะห์ข้อมูลประเภทข้อความซึ่งเป็นภาษาที่มนุษย์ใช้กันแบบธรรมชาติ กระบวนการที่สำคัญในการประมวลผลภาษาธรรมชาติ คือการทำเหมืองข้อความ (Text Mining) Vijayarani & Janani (2016) ซึ่งเป็นการขุดข้อความระบุข้อเท็จจริง ความสัมพันธ์ และคำยืนยันที่จะถูกฝังอยู่ในมวลของข้อมูลขนาดใหญ่ที่เป็นข้อความ เมื่อดึงออกมาแล้วข้อมูลนี้จะถูกแปลงเป็นรูปแบบที่มีโครงสร้างที่สามารถนำไปวิเคราะห์เพิ่มเติม การทำเหมืองข้อความใช้วิธีการที่หลากหลายในการประมวลผลข้อความ ซึ่งเป็นหนึ่งในวิธีที่สำคัญที่สุด ขั้นตอนการทำเหมืองข้อความ มีขั้นตอนพื้นฐานที่เกี่ยวข้องในการเตรียมเอกสารข้อความที่ไม่มีโครงสร้างสำหรับการวิเคราะห์เชิงลึก ได้แก่ 1. การเตรียมข้อมูลข้อความ (Text Pre-processing) 2. การแปลงข้อความ (Text Transformation) 3. การเลือกลักษณะเด่น (Feature Selection) 4. การสร้างแบบจำลองการเรียนรู้ของเครื่อง (Machine Learning Modeling) 5. การประเมิน (Evaluation) และ 6. การนำไปประยุกต์ใช้ (Deployment)

ขบวนการสำคัญอีกประการหนึ่งในการเตรียมข้อมูลสำหรับการประมวลผลภาษาธรรมชาติคือการเก็บรวบรวมข้อมูล แหล่งข้อมูลมีมากมายหลายแหล่ง ทั้งที่เป็นข้อมูลที่จัดเก็บในองค์แบบในรูปแบบไฟล์อิเล็กทรอนิกส์ชนิดต่าง ๆ ข้อมูลที่จัดเก็บในฐานข้อมูลแบบมีโครงสร้างและแบบไม่มีโครงสร้าง ข้อมูลที่อยู่บนเว็บไซต์ทั่วไปและบนสื่อสังคมออนไลน์ วิธีการสืบค้นสารสนเทศ (Information Retrieval: IR) (Krallinger et al., 2017) เป็นเทคนิควิธีสำหรับการสืบค้นข้อมูล ซึ่งสามารถนำมาประยุกต์ใช้ในการเก็บรวบรวมข้อมูล สำหรับข้อมูลที่อยู่บนเว็บไซต์ทั่วไป เทคนิคการสืบค้นจะเป็นการใช้เว็บครอว์เลอร์ค้นหาข้อมูลและสกัดข้อมูลเฉพาะที่ต้องใช้งาน ในกรณีการสร้างระบบตรวจสอบข่าวปลอม การสืบค้นสารสนเทศจะเป็นการสืบค้นข่าวที่เกี่ยวข้องกับข่าวสืบค้น และเว็บครอว์เลอร์จะไปสืบค้นและรับข้อความข่าวที่มีความคล้ายกลับข่าวสืบค้น รวมทั้งข้อมูลเมตาที่เกี่ยวข้องกับแหล่งข่าวเพื่อใช้ในการจำแนกชนิดข่าว ข่าวสืบค้นที่ได้รับกลับมามีจะถูกส่งต่อไปยังการประมวลผลภาษาธรรมชาติ เพื่อสกัดลักษณะเด่นของข่าวก่อนส่งต่อไปยังการเรียนรู้ของเครื่องในขั้นตอนถัดไป

การเรียนรู้ของเครื่อง (Machine Learning: ML) (Gori, 2018) เป็นเซตย่อยของปัญญาประดิษฐ์ ที่มีความสามารถในการเรียนรู้ข้อมูลสอนที่จัดเตรียมไว้ล่วงหน้า การเรียนรู้ของเครื่องแบ่งตามชนิดการสอนได้หลายกลุ่มหลัก ๆ ได้แก่ การเรียนรู้แบบมีผู้สอน (Supervised Learning) การเรียนรู้แบบไม่มีผู้สอน (Unsupervised Learning) การเรียนรู้แบบกึ่งมีผู้สอน (Semi-supervised Learning) และการเรียนรู้แบบเสริมแรง (Reinforcement Learning) สำหรับการเรียนรู้แบบมีผู้สอนยังสามารถแบ่งการทำงานได้อีกสองแบบ ได้แก่ การถดถอย (Regression)

และการจำแนกประเภทหรือกลุ่ม (Classification) สำหรับการสร้างแบบจำลองตรวจจำวปลอมจะอยู่ในกลุ่มการจำแนกประเภทซึ่งมีอยู่มากมายหลายเทคนิค ในที่นี้จะขอลำดับดังต่อไปนี้ การถดถอยโลจิสติกส์ เทคนิคเพื่อนบ้านใกล้เคียง โครงข่ายประสาทเทียมชนิดเพอร์เซ็ปตรอนหลายชั้น ชัฟฟอร์ตเวกเตอร์แมชชีน นาอ์ฟเบย์ ต้นไม้ตัดสินใจ ป่าสุ่ม และโครงข่ายประสาทเทียมชนิดความจำระยะสั้นแบบยาว

การถดถอยโลจิสติกส์ (Logistics Regression: LR) (Kleinbaum & Klein, 2010) การถดถอยโลจิสติกส์เป็นแบบจำลองทางสถิติที่ในรูปแบบพื้นฐานใช้ฟังก์ชันโลจิสติกส์เพื่อสร้างแบบจำลองตัวแปรตามแบบไบนารี การถดถอยโลจิสติกส์เป็นการประมาณค่าพารามิเตอร์ของแบบจำลองโลจิสติกส์ การถดถอยโลจิสติกส์สามารถนำไปใช้การวิเคราะห์เชิงคาดการณ์ตัวแปรแบบไบนารี เพื่ออธิบายข้อมูลและอธิบายความสัมพันธ์ระหว่างตัวแปรอิสระอย่างน้อยหนึ่งตัวแปรกับตัวแปรตามที่เป็นค่าเอาต์พุตซึ่งมีค่าเป็นไปได้อีกสองค่าคือ “0” และ “1” ตัวแปรอิสระหรืออินพุตซึ่งอาจจะเป็นค่าระดับ ค่าลำดับ ค่าเป็นช่วง อัตราส่วน ไบนารี หรือค่าต่อเนื่อง อย่างน้อยหนึ่งตัวแปร ในการตัดสินใจเอาต์พุตเป็นการรวมกันแบบเชิงเส้นของตัวแปรอิสระ ฟังก์ชันแบบโลจิสติกส์ใช้แปลงการรวมกันแบบเชิงเส้นของอินพุตเป็นความน่าจะเป็นซึ่งเป็นช่วงระหว่าง 0 ถึง 1 ที่สอดคล้องกับค่าเป้าหมายสองกลุ่ม “0” หรือ “1”

วิธีการเพื่อนบ้านใกล้เคียง (K-Nearest Neighbor: KNN) (Guo et al., 2003) เป็นการเรียนรู้ของเครื่องแบบมีผู้สอน และเป็นวิธีการที่ได้รับความนิยมในการใช้งานอย่างมาก เนื่องจากความง่ายในการประยุกต์ใช้งานแต่มีประสิทธิภาพสูง และสามารถนำไปประยุกต์ใช้กับงานได้อย่างหลากหลาย โดยเฉพาะอย่างยิ่งด้านการจำแนกข้อมูล ขั้นตอนการทำงานของวิธีการเพื่อนบ้านใกล้เคียง ได้แก่ 1. เก็บข้อมูลชุดสอนเป็นแบบจำลอง 2. กำหนดค่า K เป็นจำนวนสมาชิกเพื่อนบ้านใกล้เคียง 3. คำนวณหาระยะห่างระหว่างจุดด้วยวิธีต่าง ๆ อาจจะใช้ Euclidian Distance, Cosine Distance หรือวิธีอื่น ๆ เป็นการวัดค่าระยะห่างระหว่างข้อมูลทดสอบกับข้อมูลชุดสอน 4. เรียงลำดับระยะห่างระหว่างข้อมูลทดสอบและข้อมูลสอน โดยพิจารณาจากข้อมูลจำนวน K เรคคอร์ดที่ใกล้ที่สุด และ 5. ตัดสินใจโดยการโหวตผลลัพธ์จากค่าเป้าหมายของข้อมูลชุดสอนที่อยู่ใกล้ชุดทดสอบที่สุดจาก K เรคคอร์ด

โครงข่ายประสาทเทียมชนิดเพอร์เซ็ปตรอนหลายชั้น (Multilayer Perceptron: MLP) (Hagan et al., 2014) มีพื้นฐานมาจากการจำลองการทำงานของสมองมนุษย์ ด้วยโปรแกรมคอมพิวเตอร์ จุดมุ่งหมายของโครงข่ายประสาทเทียมคือต้องการให้คอมพิวเตอร์มีความชาญฉลาดในการเรียนรู้เหมือนที่มนุษย์มีการเรียนรู้ สามารถฝึกฝนได้ และสามารถนำความรู้และทักษะ รวมทั้งสามารถนำไปประยุกต์ใช้ได้ดีกับปัญหาด้านการจำแนกกลุ่ม การถดถอย และการจัดกลุ่ม เทคนิคนี้มักถูกเรียกว่า “Black Box” เนื่องจากการทำงานมีความซับซ้อนมากกว่าเทคนิคอื่นๆ ก่อนข้างมาก มีการนำมาประยุกต์ใช้ในงานด้านต่าง ๆ ข้อดีของการจำแนกกลุ่มด้วยวิธีโครงข่ายประสาทเทียม ได้แก่ ความมีประสิทธิภาพในการพยากรณ์ข้อมูลสูง และรองรับการจำแนกข้อมูลได้ทั้งข้อมูลที่เป็นลักษณะตัวเลขต่อเนื่อง ข้อจำกัดของการจำแนกกลุ่มด้วยแบบจำลองโครงข่ายประสาทเทียม ได้แก่ 1) บางครั้งอาจเกิดเหตุการณ์เรียนรู้ข้อมูลสอนมากเกินไปทำให้ แบบจำลองเรียนรู้ได้ผลลัพธ์ที่ดีเฉพาะในข้อมูลสอน แต่กลับมีประสิทธิภาพไม่ดีเมื่อนำมาใช้กับข้อมูลในสภาพแวดล้อมจริง 2) กระบวนการทำงานของแบบจำลองเป็นแบบ Black block จึงไม่

สามารถทราบเหตุผลการตัดสินใจของแบบจำลอง 3) ใช้เวลาในการเรียนรู้และใช้หน่วยความจำมาก หากข้อมูลมีขนาดใหญ่ และ 4) ไม่คงทนต่อข้อมูลรบกวนและข้อมูลแปลกปลอม

ซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) (Cortes & Vapnik, 1995) เป็นเครื่องมือสำหรับการเรียนรู้แบบใหม่บนพื้นฐานของการเรียนรู้ทฤษฎีทางสถิติ มีประโยชน์อย่างมากในการนำไปใช้จำแนกข้อมูลหรือการรู้จำรูปแบบของข้อมูลโดยวิธี SVM ได้มีงานวิจัยหลายงานที่นำวิธีการดังกล่าวไปประยุกต์ใช้ SVM มีแนวคิดของทฤษฎีดังนี้ 1) การลดความเสี่ยงเชิงโครงสร้างให้ต่ำที่สุด เป็นแนวคิดที่แสดงถึงขอบเขตของความเสียหายหรือความน่าจะเป็นที่จะเกิดความผิดพลาดของการเรียนรู้ มีประโยชน์สำหรับการนำมาใช้ในกระบวนการเรียนรู้เพื่อหาฟังก์ชันในการตัดสินใจเพื่อให้เกิดความผิดพลาดน้อยที่สุด 2) ปริภูมิแสดงลักษณะสำคัญกับส่วนของเคอร์เนลฟังก์ชัน ที่จะทำการแมปข้อมูลจากปริภูมินำเข้าไปสู่ปริภูมิแสดงลักษณะที่สำคัญ เพื่อประโยชน์ในการสร้างฟังก์ชันสำหรับการตัดสินใจที่มีลักษณะไม่เป็นไปในรูปแบบเชิงเส้นกับข้อมูลในปริภูมินำเข้า ทำให้การแบ่งกลุ่มข้อมูลมีความยืดหยุ่นมากขึ้น และ 3) ระบายหลายมิติที่ใช้แยกข้อมูลได้ดีที่สุด เป็นแนวคิดหลักของเทคนิค SVM เพื่อทำการหาระนาบที่มีระยะห่างจากข้อมูลทำการแบ่งให้มากที่สุด สำหรับการแบ่งกลุ่มข้อมูลสองกลุ่มออกจากกัน งานวิจัยทั่วไปพบว่า SVM มีข้อดีด้านความคงทนต่อข้อมูลรบกวน

นาอิวเบย์ (Naïve Bayesian: NB) (Yager, 2006) เป็นตัวแยกประเภทในกลุ่มความน่าจะเป็น ประยุกต์ใช้ทฤษฎีของเบย์ (Baye's Theorem) NB มีสมมติฐานความเป็นอิสระที่ระหว่างตัวแปรต่าง ๆ นาอิวเบย์ได้รับการศึกษาอย่างกว้างขวางตั้งแต่ทศวรรษ 1960 ยังคงเป็นวิธีที่นิยมถึงปัจจุบัน สำหรับการจัดหมวดหมู่ข้อความ การจำแนกเอกสาร เช่น การจำแนกอีเมล กฎหมาย กีฬา และการเมือง เป็นต้น โดยนาอิวเบย์จะใช้ข้อมูลสอนที่เป็นค่าจากข้อความแปลงเป็นค่าคุณลักษณะที่ถูกประมวลผลไว้ล่วงหน้า ด้วยความไม่ซับซ้อนของนาอิวเบย์จึงได้รับความสนใจถูกนำไปใช้ในแอปพลิเคชันมากมาย นาอิวเบย์สามารถปรับขนาดของงานขนาดใหญ่ โดยต้องการพารามิเตอร์จำนวนหนึ่งเป็นเชิงเส้นในจำนวนตัวแปรอินพุตและเอาต์พุต นาอิวเบย์เป็นเทคนิคที่นิยมใช้เปรียบเทียบกับเทคนิคอื่น ๆ ถือเป็นบรรทัดฐาน

วิธีต้นไม้ตัดสินใจ (Decision Tree: DT) (Myles et al., 2004) มีลักษณะโครงสร้างแบบต้นไม้หัวกลับ เป็นการนำข้อมูลมาสร้างแบบจำลองการพยากรณ์ในรูปของต้นไม้ตัดสินใจ โดยมีการใช้แอตทริบิวต์ของข้อมูลกำหนดเป็นโหนดต่าง ๆ และกิ่งก้านของต้นไม้เป็นค่าของข้อมูลในแต่ละแอตทริบิวต์ ต้นไม้ตัดสินใจเริ่มต้นด้วยโหนดราก แยกข้อมูลตามเงื่อนไขไปตามกิ่งก้าน เชื่อมต่อกับโหนดระหว่างกลาง และแยกข้อมูลตามกิ่งก้านตามเงื่อนไขต่อไปเรื่อย ๆ ไปสิ้นสุดที่โหนดใบซึ่งเป็นโหนดตัดสินใจ ซึ่งต้นไม้ตัดสินใจนั้นทำงานแบบเรียนรู้แบบมีผู้สอน สามารถประยุกต์ใช้ในการจำแนกประเภทได้ จากกลุ่มตัวอย่างของข้อมูลที่ได้กำหนดได้ก่อนล่วงหน้าซึ่งใช้เป็นข้อมูลสอน และสามารถพยากรณ์กลุ่มของรายการที่ยังไม่เคยพบมาก่อนได้ ต้นไม้ตัดสินใจเป็นแบบจำลองที่โครงสร้างไม่ซับซ้อน สร้างได้ง่าย และตีความเป็นกฎการตัดสินใจได้ จึงเป็นที่นิยมใช้ในกระบวนการทำเหมืองข้อมูลและนำไปใช้ในการตัดสินใจ

ป่าสุ่ม (Random Forest: RF) (Svetnik et al., 2003) หรือป่าตัดสินใจสุ่ม (Random Decision Forest) เป็นวิธีการเรียนรู้ทั้งมวล (Ensemble Learning) ถูกนำไปใช้สำหรับการจำแนกประเภท การถดถอย และงานอื่น ๆ มากมาย การสร้างแบบจำลองป่าสุ่มเป็นการสร้างต้นไม้การตัดสินใจจำนวนมากจากข้อมูลสอน โดยมีแนวคิดว่าการช่วยกันคิดตัดสินใจมีโอกาสลดข้อผิดพลาดได้ ในการจำแนกประเภทผลลัพธ์ในการตัดสินใจของป่าสุ่มได้จากประเภทจากการตัดสินใจโดยต้นไม้ส่วนใหญ่ สำหรับงานการถดถอยใช้ค่าเฉลี่ยหรือการคาดการณ์ของต้นไม้แต่ละต้นในการให้คำตอบตัดสินใจ โดยทั่วไปแล้วป่าสุ่มจะมีประสิทธิภาพดีกว่าต้นไม้ตัดสินใจเดี่ยว โดยลักษณะข้อมูลอาจส่งผลต่อประสิทธิภาพการทำงานได้ด้วย

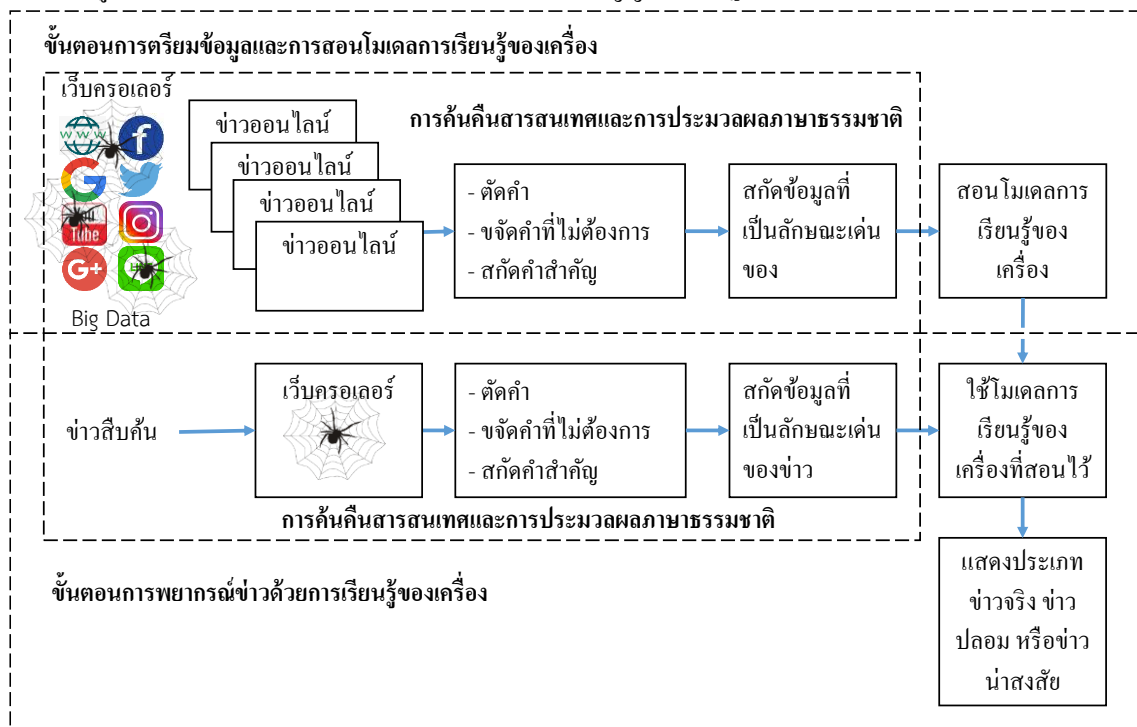
โครงข่ายประสาทเทียมแบบความจำระยะสั้นระยะยาว (Long Short-Term Memory Neural Network: LSTM) (Hochreiter & Schmidhuber, 1997) เป็นการเรียนรู้เชิงลึกกลุ่มเดียวกับโครงข่ายประสาทเทียมแบบวนซ้ำ (RNN) LSTM มีการย้อนกลับของข้อมูล จึงทำให้ไม่เพียงแต่สามารถประมวลผลข้อมูลได้แต่ยังสามารถประมวลผลข้อมูลแบบลำดับได้เป็นอย่างดี LSTM สามารถใช้ได้กับงานต่าง ๆ เช่น การจำแนกความรู้สึกจากเอกสาร การรู้จำลายมือที่เชื่อมต่อ การรู้จำเสียงพูด และการตรวจจับความผิดปกติของข้อมูลเครือข่าย นิวรอนของ LSTM ทั่วไปเรียกว่าเซลล์ ประกอบด้วยประตูเข้า ประตูส่งออก และประตูลืม เซลล์จะจดจำค่าในช่วงเวลาต่าง ๆ และประตูทั้งสามจะควบคุมการไหลของข้อมูลเข้าและออกจากเซลล์ LSTM เหมาะสมอย่างยิ่งกับการจำแนกการประมวลผลและการทำนายตามข้อมูลอนุกรมเวลา เนื่องจากอาจมีช่วงเวลาที่ไม่มีทราบระยะเวลาระหว่างเหตุการณ์สำคัญในอนุกรมเวลา LSTM ถูกการพัฒนาขึ้นเพื่อแก้ปัญหาการลู่สู่ค่าอนันต์ และการลู่เข้าสู่ศูนย์ของค่าความลาดชันค่าผิดพลาดที่พบในขณะสอนของ RNN แบบเดิม

การดำเนินงานวิจัย

การดำเนินงานวิจัย ประกอบด้วย 4 ขั้นตอนหลัก ได้แก่ 1) การออกแบบกรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์ 2) การนำเสนอลักษณะเด่นจากการวิเคราะห์ภาษาธรรมชาติ 3) การสร้างแบบจำลองการเรียนรู้ของเครื่องสำหรับการตรวจข่าวปลอมภาษาไทย และ 4) การพัฒนาเว็บแอปพลิเคชันตรวจข่าวปลอมภาษาไทย

เนื้อหาข่าวปลอมเกิดขึ้นจากมีผู้ที่ตั้งใจให้มีความหมายเปลี่ยนไปจากข้อมูลเดิมโดยมีวัตถุประสงค์แตกต่างกันไป ข่าวปลอมจึงเป็นข้อมูลที่กว้างและมีความเป็นพลวัตสูง ทำให้ยากต่อการสร้างการเรียนรู้ของเครื่องให้เป็นระบบตรวจข่าวปลอมที่สามารถเรียนรู้ตามเนื้อหาที่เปลี่ยนไป อย่างไรก็ตามมีความเป็นไปได้ที่จะสร้างระบบตรวจข่าวปลอมมีประสิทธิภาพสูง จากการศึกษาพบว่าผู้ผลิตข่าวนิยมเผยแพร่ข่าวออนไลน์เพื่อผู้คนสามารถเข้าถึงข่าวได้อย่างรวดเร็วผ่านอินเทอร์เน็ตจากทุกที่ อีกทั้งสื่อสังคมออนไลน์ได้เก็บทั้งข่าวจริงและข่าวปลอมบนคลาวด์เซิร์ฟเวอร์ มีผู้อ่านข่าวได้แสดงความคิดเห็นในเว็บโซเชียลมีเดียจำนวนมาก จึงเป็นแหล่งข้อมูลที่สำคัญในการตรวจข่าวได้แบบอัตโนมัติและมีประสิทธิภาพสูงได้

ในงานวิจัยนี้ผู้วิจัยเสนอกรอบงานการตรวจจับข่าวปลอมด้วยการเรียนรู้ด้วยเทคนิคปัญญาประดิษฐ์ ดังรูปที่ 1 โดยกรอบงานตรวจสอบข่าวปลอมที่นำเสนอ แบ่งเป็นสองขั้นตอน คือ ขั้นตอนการเตรียมข้อมูลและการสอนแบบจำลองการเรียนรู้ของเครื่อง และขั้นตอนการพยากรณ์ข่าวด้วยการเรียนรู้ของเครื่อง ทั้งสองขั้นตอนมีความเกี่ยวข้องกับโมดูลสามโมดูลหลัก ได้แก่ การค้นคืนสารสนเทศ (IR) การประมวลผลภาษาธรรมชาติ (NLP) และการเรียนรู้ของเครื่อง (ML) โมดูล IR ดึงข้อมูลข่าวสารจากอินเทอร์เน็ตตามข่าวสืบค้นที่ป้อนโดยผู้ใช้ ผลลัพธ์จากการค้นคืนเป็นเนื้อหาข่าวที่เกี่ยวข้องจากแหล่งข่าวออนไลน์มากมาย โมดูล NLP จะวิเคราะห์เอกสารที่ได้รับกลับคืนมา โดยดำเนินการตัดคำ ขจัดคำที่ไม่ต้องการ และสกัดคุณลักษณะข่าว โมดูล ML แบ่งประเภทข่าวออกเป็นสามประเภท ได้แก่ ข่าวจริง ข่าวปลอม และข่าวน่าสงสัย และในส่วนของ “การแสดงผลข้อมูลข่าวที่เกี่ยวข้องกับข่าวสืบค้น” รูปที่ 1 แสดงกรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์



รูปที่ 1 กรอบการตรวจข่าวปลอมด้วยเทคนิคปัญญาประดิษฐ์

'ข่าวไม่ปลอม', 'ข่าวจริง', 'ข่าวไม่เท็จ', 'ข่าวไม่บิดเบือน', 'ไม่ปลอม', 'เป็นข่าวจริง', 'เป็นเรื่องจริง', 'ไม่เท็จ', 'ไม่บิดเบือน', 'เรื่องจริง', 'ข้อความจริง', 'ข้อมูลจริง', 'เป็นข้อมูลจริง', 'ข้อมูลถูก', 'แชร์ต่อได้', 'เป็นความจริง', 'มีมูล'

รูปที่ 2 ตัวอย่างคำที่มักปรากฏในข่าวจริงที่แตกต่างจากคำในข่าวปลอม

'การติดต่อ', 'การหมิ่นประมาท', 'การหลอกลวง', 'การเล่นตลก', 'การโฆษณาชวนเชื่อ', 'ข่าวขยะ', 'ข่าวที่ผิดพลาด', 'ข่าวที่ไม่ถูกต้อง', 'ข่าวที่ไม่น่าเชื่อถือ', 'ข่าวบิดเบือน', 'ข่าวปลอม', 'ข่าวปลอม', 'ข่าวผิด', 'ข่าวมั่ว', 'ข่าวลบ', 'ข่าวเท็จ', 'ข่าวไม่จริง', 'ข่าวไม่ถูกต้อง', 'ข้อความข่าวปลอม', 'ข้อความบิดเบือน', 'ข้อความปลอม', 'ข้อความเท็จ', 'ข้อมูลปลอม', 'ข้อมูลผิด', 'ข้อมูลผิดพลาด', 'ข้อมูลเท็จ', 'ข้อมูลไม่จริง', 'ข้อเท็จจริงเท็จ', 'ข้อเท็จจริงไม่ถูกต้อง', 'ข้อเท็จจริงไร้ค่า', 'ความเชื่อผิดๆ', 'ความเท็จ', 'คำกล่าวอ้าง', 'คำพูดเหลวไหล', 'คำลวง', 'คำเท็จ', 'คำโกหก', 'จับโป๊ะ', 'ชวนเชื่อ', 'ต่อแหล', 'บิดเบือน', 'ปลอม', 'ปลอม', 'มั่ว', 'มโน', 'รายงานเท็จ', 'สื่อสนั่น', 'สื่อหึ่ง', 'ข่าวหลอกลวง', 'อวดอ้าง', 'เท็จ', 'เป็นข้อมูลเท็จ', 'เป็นเฟคนิส', 'เรื่องราวปลอม', 'เรื่องเท็จ', 'เหตุการณ์หลอก', 'โฆษณาชวนเชื่อ', 'โยงมั่ว', 'ไม่น่าเชื่อถือ', 'ไม่มีมูลความจริง', 'ไม่มีอยู่จริง', 'ไม่เป็นความจริง', 'ไม่ใช่ข่าวจริง', 'ไม่ใช่เรื่องจริง'

รูปที่ 3 ตัวอย่างคำที่มักปรากฏในข่าวปลอมที่แตกต่างจากคำในข่าวจริง

ผู้วิจัยทำการเก็บข้อมูลด้วยวิธีการค้นคืนสารสนเทศด้วยเว็บครอว์เลอร์เพื่อเก็บข้อมูลตั้งแต่วันที่ 1 สิงหาคม 2563 ถึงวันที่ 31 พฤษภาคม 2564 ข้อมูลที่ได้รับคืนกลับมาจะใช้สำหรับการสอนแบบจำลองการเรียนรู้ของเครื่อง จำนวน 44,271 ข้อความข่าว แบ่งเป็น 3 ประเภท ได้แก่ ข่าวจริง (real) ข่าวปลอม (fake) และข่าวน่าสงสัย (suspicious) การระบุประเภทเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย นั้น ผู้วิจัยนำข่าวที่ได้รับจากค้นคืนช่วงเริ่มต้นในการเก็บ มาวิเคราะห์เพื่อหาองค์ประกอบของข่าวทั้งสามประเภท เพื่อหาฟีเจอร์หรือลักษณะเด่นของข่าว 5 ค่า ได้แก่ จำนวนคำที่เป็นลักษณะข่าวปลอม (score_fake) จำนวนคำที่เป็นลักษณะข่าวจริง (score_real) ค่าความคล้ายของข่าว (sim_matched) จำนวนโดเมนที่พบระบุข่าวปลอม (domain_fake) จำนวนโดเมนที่พบระบุข่าวจริง (domain_real) สำหรับการคำนวณค่าลักษณะเด่นที่เป็นค่าความคล้ายของข่าว (sim_matched) วัดจากค่าความเหมือนแบบโคไซน์ (cosine similarity) (Soyusiwaty & Zakaria, 2018) ผู้เขียนเป็นสคริปต์เป็นกฎเพื่อจัดเก็บข้อมูลข่าวตามฟีเจอร์และทำการระบุค่าเป้าหมายประเภทของข่าวโดยอัตโนมัติ ซึ่งระบุประเภทเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย จากนั้นผู้วิจัยตรวจสอบความถูกต้องอีกครั้งเพื่อพร้อมใช้สำหรับการสอนให้เครื่องเรียนรู้ รูปที่ 2 แสดงตัวอย่างคำที่มักปรากฏในข่าวจริงที่แตกต่างจากคำในข่าวปลอม และรูปที่ 3 แสดงตัวอย่างคำที่มักปรากฏในข่าวปลอมที่แตกต่างจากคำในข่าวจริง

ผู้วิจัยทำให้สมดุลข้อมูลในแต่ละกลุ่มโดยวิธีการสุ่มเพิ่มแบบ SMOTE (Chawla *et al.*, 2002) เป็นกลุ่มละ 17,740 ข้อความ รวมเป็นจำนวน 53,220 ข้อความข่าว ดังแสดงในตารางที่ 1

ตารางที่ 1 ข้อมูลที่จัดเก็บด้วยวิธีการค้นคืนสารสนเทศด้วยเว็บครอว์เลอร์

ชนิดข่าว	จำนวนข่าว	จำนวนสมมูล (คู่เพิ่ม)
ปลอม (Fake)	17,740	17,740
จริง (Real)	13,072	17,740
น่าสงสัย (Suspicious)	13,459	17,740
รวม	44,271	53,220

งานวิจัยนี้ข้อมูลข่าวสืบค้นถูกสกัดเป็นลักษณะเด่นของข้อมูลจำนวน 5 ค่า ได้แก่ score_fake, score_real, sim_matched, domain_fake และ domain_real ค่าเป้าหมายประกอบด้วยข่าวปลอม (fake) ข่าวจริง (real) และน่าสงสัย (suspicious) ข่าวสืบค้นจะถูกนำไปเทียบกับข่าวออนไลน์ที่มีความคล้ายคลึงกัน ผู้วิจัยใช้หลักการการประมวลผลภาษาธรรมชาติเพื่อสกัดลักษณะเด่น 5 ค่า วิธีการสกัดลักษณะเด่น เริ่มจากการผลการค้นคืนสารสนเทศที่เป็นข่าวคล้ายคลึงกับข่าวสืบค้น ทำการตัดค่าและนับจำนวนค่าที่เป็นลักษณะข่าวปลอม (score_fake) นับจำนวนค่าที่เป็นลักษณะข่าวจริง (score_real) คำนวณค่าความคล้ายของข่าว (sim_matched) นับจำนวนโดเมนที่พบระบุข่าวปลอม (domain_fake) นับจำนวนโดเมนที่พบระบุข่าวจริง (domain_real) ตัวอย่างข้อมูลข่าวสืบค้นและลักษณะเด่นแสดงดังตารางที่ 2

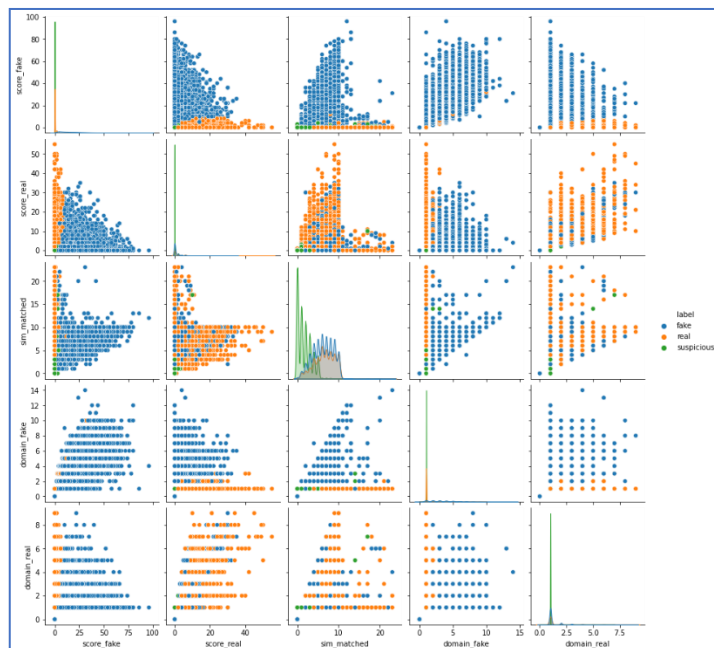
ตารางที่ 2 ตัวอย่างข้อมูลข่าวสืบค้นและลักษณะเด่น

ข่าวสืบค้น	ประเภท	score_fake	score_real	sim_matched	domain_fake	domain_real
กรม. เกาวันที่ 4 และ 7 ก.ย. 63 เป็นวันหยุดชดเชยวันสงกรานต์	real	2	4	21	1	3
ฆ่าเชื้อโควิด-19 ด้วยการกินผลไม้ ที่มีฤทธิ์เป็นด่าง	fake	7	0	12	6	1
ขุนพลเคนมาร์กกับการสร้าง กำแพงปกป้องเอริเคนในวันที่ หมดสติบนสนามของศึกยูโร 2020	suspicious	0	0	0	0	0
หมู-ไก่ เป็นโรคเอดส์	fake	31	4	23	14	4
กรมอนามัย รุกโรงงานขนาด ใหญ่ ประเมิน Thai Stop COVID Plus ครบ 100 % 15 มิถุนายนนี้	suspicious	0	0	2	0	0
รพท. เปิดมหรหรรษ์ใกล้เกลี่ย สินเชื่อเช่าซื้อรถ ผ่านทาง ออนไลน์	real	4	6	9	1	2

ตารางที่ 3 ค่าความสัมพันธ์ของตัวแปรในข้อมูลสำหรับการสอนและทดสอบแบบจำลอง

	score_fake	score_real	sim_matched	domain_fake	domain_real	fake	Real	suspicious
score_fake	1.000	0.081	0.424	0.901	0.103	0.724	-0.377	-0.398
score_real	0.081	1.000	0.201	0.097	0.785	0.042	0.223	-0.266
sim_matched	0.424	0.201	1.000	0.440	0.236	0.360	0.303	-0.683
domain_fake	0.901	0.097	0.440	1.000	0.121	0.693	-0.363	-0.378
domain_real	0.103	0.785	0.236	0.121	1.000	0.056	0.129	-0.187
Fake	0.724	0.042	0.360	0.693	0.056	1.000	-0.529	-0.540
Real	-0.377	0.223	0.303	-0.363	0.129	-0.529	1.000	-0.428
suspicious	-0.398	-0.266	-0.683	-0.378	-0.187	-0.540	-0.428	1.000

ตารางที่ 3 แสดงค่าลักษณะเด่นที่แยกออกมาที่มีความสัมพันธ์กับเป้าหมายโดยลักษณะเด่น sim_matched มีค่าสหสัมพันธ์เชิงบวกกับทั้งข่าวปลอมและข่าวจริง ลักษณะเด่น score_fake และ domain_fake มีค่าสหสัมพันธ์เท่ากับ 0.724 และ 0.693 ตามลำดับ ซึ่งแสดงถึงความมีผลต่อการคาดการณ์ข่าวปลอม นอกจากนี้ลักษณะเด่น score_real และ domain_real มีค่าสหสัมพันธ์เชิงบวกกับข่าวจริง โดยเท่ากับ 0.223 และ 0.129 ตามลำดับ ในภาพรวมข้อมูลที่สกัดลักษณะเด่นจากงานวิจัยนี้สามารถนำไปสร้างการเรียนรู้ด้วยเครื่องเพื่อทำนายข่าวปลอมและข่าวจริงได้



รูปที่ 3 การพล็อตแบบกระจายแสดงความสัมพันธ์ของข้อมูล

เพื่อการมองข้อมูลเป็นภาพและวิเคราะห์ความสัมพันธ์ ผู้วิจัยใช้วิธีการแสดงเป็นการพล็อตแบบกระจายรูปที่ 3 แสดงการพล็อตข้อมูลแบบกระจาย เป็นการแสดงตำแหน่งของข้อมูลตามความสัมพันธ์ระหว่าง 2 ตัวแปร โดยหนึ่งตัวแปรในแนวนอนแต่ละแถว และแนวตั้งเป็นอีกหนึ่งตัวแปรในแต่ละคอลัมน์ ในแต่ละภาพย่อยเป็นความสัมพันธ์ของ 2 ตัวแปร การพล็อตแบบกระจายทำให้พบข้อสังเกตได้ว่าข้อมูลสามารถจำแนกได้ง่ายในหลาย ๆ ภาพที่เป็นคู่ความสัมพันธ์ระหว่างลักษณะเด่น เป็นการยืนยันว่าข้อมูลที่ผ่านการประมวลผลสามารถนำไปสร้างการเรียนรู้ของเครื่องได้

ในการสร้างแบบจำลองการเรียนรู้ของเครื่องเพื่อจำแนกข่าวจริงและข่าวปลอม จะต้องจัดเตรียมข้อมูลจำนวนมากเพียงพอสำหรับสอนแบบจำลองการเรียนรู้ของเครื่อง งานวิจัยนี้ผู้วิจัยใช้ข้อมูลประเภทข้อความข่าวก่อนการนำไปสอนให้กับตัวแบบการเรียนรู้ของเครื่องต้องทำการประมวลผลคำด้วยเทคนิควิธีการประมวลผลภาษาธรรมชาติ ส่วนอัลกอริทึมการเรียนรู้ของเครื่องที่เลือกใช้มีหลายเทคนิค เช่น LR, KNN, MLP, SVM, NB, DT, RF และ LSTM เทคนิคเหล่านี้เป็นเทคนิคการเรียนรู้ของเครื่องชนิดมีผู้สอน โดยมีรายละเอียดดังนี้ 1) การเตรียมข้อมูล ดำเนินการหลังจากขั้นตอนการรวบรวมข้อมูลด้วยการสืบค้นสารสนเทศในเว็บด้วยเว็บเบราว์เซอร์ จากนั้นนำข้อมูลที่ได้อิงประมวลผลภาษาธรรมชาติ ได้คุณลักษณะเด่นอยู่ในรูปแบบที่สามารถวิเคราะห์ได้ด้วยเทคนิคการเรียนรู้ของเครื่อง การประมวลผลภาษาธรรมชาติ 2) การสร้างแบบจำลองหรือตัวแบบการเรียนรู้ของเครื่อง โดยทำการนำข้อมูลที่ผ่านการเตรียมข้อมูลแล้ว เข้าสู่การสร้างแบบจำลองโดยใช้เทคนิคการเรียนรู้แบบมีผู้สอน โดยการสอนแบบจำลองใช้ข้อมูล 53,220 เรคคอร์ด สอนด้วยวิธีการทดสอบแบบไขว้ 10 ชุดข้อมูล (10-Fold Cross Validate) 3) การทดสอบและประเมินประสิทธิภาพแบบจำลอง โดยนำชุดข้อมูลที่เครื่องคอมพิวเตอร์ยังไม่เคยเรียนรู้มาทำการทดสอบแบบจำลองว่ามีค่าประสิทธิภาพในการวิเคราะห์ โดยประเมินจาก ค่าความถูกต้อง (Accuracy) ความแม่นยำ (Precision) ค่าการเรียกคืน (Recall) และค่าถ่วงดุล (F-Measure) และ 4) การนำแบบจำลองที่ดีที่สุดไปประยุกต์ใช้งาน โดยการพัฒนาเป็นเว็บแอปพลิเคชันระบบตรวจข่าวปลอม

การพัฒนาระบบตรวจข่าวปลอม ผู้วิจัยเน้นการพัฒนาให้ทำงานบนคลาวด์เพื่อรองรับความพร้อมใช้งาน เน้นการเลือกใช้ซอฟต์แวร์แบบเปิดทั้งหมด เช่น ระบบปฏิบัติการ Ubuntu 20.04 LTS ใช้ Python เวอร์ชัน 3.8 เป็นภาษาในการพัฒนาเป็นหลัก ใช้เว็บเฟรมเวิร์กเป็น Django เวอร์ชัน 3.1.5 ใช้ฐานข้อมูลเป็นแบบ MongoDB เวอร์ชัน 4.4.6 สำหรับจัดเก็บข่าว และใช้ MySQL เวอร์ชัน 8.0 สำหรับจัดเก็บข้อมูลผู้ใช้งานและเมตาดาตาของระบบ สำหรับระบบสืบค้นข้อมูลใช้เทคโนโลยีที่สำคัญ ๆ ได้แก่ Selenium, Requests, Google, Chrome, และ Firefox สำหรับเครื่องมือประมวลผลภาษาธรรมชาติ Pythainlp เวอร์ชัน 2.3.1 (Phatthiyaphaibun et al., 2016) เครื่องมือในการสร้างการเรียนรู้ด้วยเครื่องประกอบด้วย Scikit-learn และ TensorFlow ด้านความปลอดภัยของระบบ ใช้เครื่องมือเพื่อรักษาความปลอดภัยและป้องกันการบุกรุก ได้แก่ Fail2ban, Uncomplicated Firewall และ ClamAV

ผลการดำเนินการวิจัย

การสร้างแบบจำลองการเรียนรู้ของเครื่องสำหรับตรวจจับข่าวปลอม ผู้วิจัยทำการเปรียบเทียบแบบจำลองเพื่อเลือกแบบจำลองที่ดีที่สุดในระบบตรวจจับข่าวปลอม ผู้วิจัยเลือกซอฟต์แวร์แบบเปิด ได้แก่ Python, Numpy, Pandas, Scikit-Learn, Tensorflow และเครื่องมือที่จำเป็นอื่น ๆ ในการสร้างระบบตรวจจับข่าวปลอม เครื่องมือวิเคราะห์ข้อมูลและการสร้างแบบจำลองการเรียนรู้ของเครื่อง ได้แก่ LR, MLP, DT, RF, SVM, NB และ KNN ในแพ็คเกจ Scikit-learn และ LSTM ใน TensorFlow ตัวชี้วัดประสิทธิภาพที่ใช้รวมถึงค่าความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่าการเรียกคืน (Recall) และค่าถ่วงดุล (F-Measure)

จากผลการวิจัยในการสร้างแบบจำลองการเรียนรู้ของเครื่องประกอบไปด้วย สำหรับเทคนิคการเรียนรู้ด้วยเครื่องที่ใช้ในการทดลองเปรียบเทียบในการวิจัยครั้งนี้ พบว่าทุกเทคนิคมีผลการทดสอบความสามารถในการจำแนกข่าวได้ดี ซึ่งมีค่าประสิทธิภาพที่สูงกว่า 90% ในทุกเทคนิค ยกเว้น NB ซึ่งมีการจำแนกข้อมูลทดสอบถูกต้องเพียง 78% เท่านั้น เมื่อพิจารณาถึงความถูกต้องการจำแนกข้อมูลเรียงลำดับจากสูงสุดไปต่ำสุด มีลำดับดังต่อไปนี้ DT (94.2%), RF (94.2%), SVM (94.0%), LSTM (93.81%), MLP (93.6%), LR (91.9%), KNN (91.5%) และ NB (78.2%) ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยของการเรียนรู้ของเครื่องชนิดต่าง ๆ แสดงดังตารางที่ 4

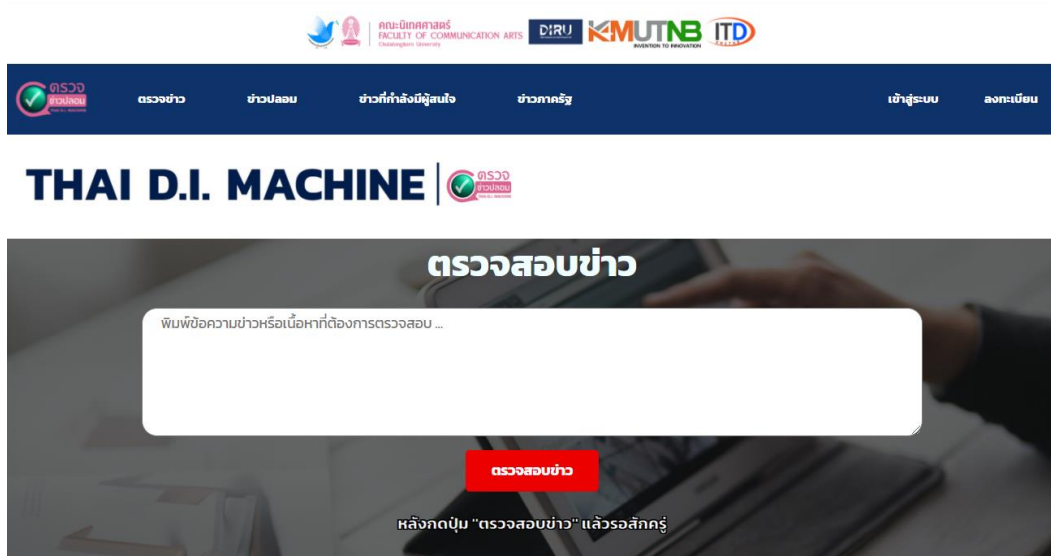
ตารางที่ 4 ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยของการเรียนรู้ของเครื่องชนิดต่าง ๆ

Metrics	LR	MLP	DT	RF	SVM	NB	KNN	LSTM
Accuracy	0.919±0.028	0.936±0.047	0.942±0.052	0.942±0.052	0.940±0.052	0.785±0.050	0.915±0.032	93.81±0.006
Precision	0.923±0.026	0.948±0.036	0.952±0.037	0.953±0.037	0.950±0.038	0.814±0.041	0.925±0.027	94.44±0.006
Recall	0.919±0.028	0.936±0.047	0.942±0.052	0.942±0.052	0.940±0.052	0.785±0.050	0.915±0.032	93.81±0.006
F-Measure	0.919±0.028	0.939±0.055	0.940±0.055	0.940±0.055	0.938±0.055	0.776±0.061	0.914±0.033	93.78±0.006

เมื่อวิเคราะห์ถึงเหตุผลในการจำแนกข้อมูลของแต่ละเทคนิค เริ่มจากเทคนิค สำหรับ 5 เทคนิคที่มีประสิทธิภาพสูงสุดด้านความถูกต้อง ตั้งแต่ 93.81% ถึง 94.2% ได้แก่ MLP, LSTM, SVM, RF และ DT ซึ่งมีความแตกต่างกันไม่มาก MLP เป็นโครงข่ายประสาทเทียมแบบหลายชั้นมีความเป็นไปได้ที่จะปรับปรุงให้มีประสิทธิภาพสูงขึ้นด้วยการเพิ่มจำนวนชั้นและพารามิเตอร์แต่จะเกิดปัญหาการจำข้อมูลสอนมากเกินไป (Overfitting) ได้เช่นกัน ส่วน LSTM เป็นเทคนิคโครงข่ายประสาทเทียมประเภทการเรียนรู้เชิงลึกซึ่งมีจำนวนพารามิเตอร์จำนวนมาก จึงมีโอกาสสูงในการปรับปรุงพารามิเตอร์ให้มีขอบเขตข้อมูลได้ดีและมีความคงเส้นคงวาทนทานต่อข้อมูลแปลกปลอมได้ดีที่สุดโดยมีส่วนเบี่ยงเบนมาตรฐานต่ำที่สุด ในขณะที่ SVM นอกจากจะมีประสิทธิภาพสูงยังมีความทนทานต่อข้อมูลรบกวนและข้อมูลแปลกปลอมได้ดีเช่นกัน สำหรับ RF ซึ่งได้ผลเทียบเท่ากับ DT แต่ RF ใช้ทรัพยากรที่มากกว่า ทั้งนี้ RF อาศัยหลักการจำแนกข้อมูลด้วยชุดต้นไม้ตัดสินใจจำนวนมาก เป็นการเพิ่มจำนวนแบบจำลองย่อยร่วมกันตัดสินใจให้มีความถูกต้องในการจำแนกได้ดียิ่งขึ้น ในส่วน KNN เป็นการจำแนกข้อมูลโดยไม่ต้องมีการสร้างแบบจำลองโดยตรงแต่จะใช้ข้อมูลชุดสอนทั้งหมดเป็นแบบจำลอง และจำแนกข้อมูลด้วยการวัดความคล้ายคลึงระหว่างข้อมูลที่จะจำแนก เทียบกับข้อมูลชุดสอน ด้วยการกำหนดค่า $K=5$ เป็นการเลือกเรคคอร์ดที่ใกล้เคียงที่สุด 5

เรคคอร์ดเพื่อโหวตผลการจำแนก KNN เองก็ถือว่าเป็นที่นิยมเพราะง่ายในการสร้าง ให้ผลลัพธ์เป็นที่ยอมรับได้เช่นกัน สำหรับ LR เป็นเทคนิคการถดถอยที่สามารถทำงานสำหรับจำแนกข้อมูลเนื่องจากใช้ฟังก์ชันลอจิสติกมอดเป็นแปลงเอาต์พุตเป็น 0 และ 1 แทนที่จะเป็นค่าต่อเนื่อง เทคนิค LR ได้ผลการจำแนกสูงกว่า 90% เช่นกันในงานวิจัยครั้งนี้ สำหรับ NB ซึ่งเป็นเทคนิคที่ใช้หลักการความน่าจะเป็น ซึ่งเป็นเทคนิคเดียวที่มีประสิทธิภาพต่ำกว่า 80% ทั้งนี้ในงานวิจัยส่วนใหญ่พบว่า NB จะเป็นเทคนิคที่เป็นรองเทคนิคอื่น ๆ จึงนิยมใช้เป็นตัวหลักในการเปรียบเทียบ ซึ่งเทคนิคใหม่ ๆ คาดหวังว่าจะต้องดีกว่า NB

เมื่อพิจารณาถึงเหตุผลของการจัดเตรียมข้อมูล ถือว่าเป็นสิ่งสำคัญก่อนป้อนข้อมูลเข้าสู่แบบจำลองการเรียนรู้ของเครื่อง หากการจัดเตรียมข้อมูลได้ไม่ดี ก็ไม่อาจสร้างแบบจำลองการเรียนรู้ของเครื่องในการจำแนกข้อมูลที่มีประสิทธิภาพสูงได้ ซึ่งในงานวิจัยครั้งนี้ใช้วิธีการสกัดลักษณะเด่นของคำที่มักปรากฏในข่าวจริง และข่าวปลอมโดยการนับความถี่ของคำเหล่านั้นเพื่อใช้เป็นข้อมูลประจำของแต่ละข่าวที่สืบค้น อีกหนึ่งลักษณะเด่นคือค่าความคล้ายของข่าวสืบค้นกับเนื้อหาในแหล่งข่าวที่เป็นปัจจัยหลักทำให้เกิดอำนาจจำแนกข้อมูลได้ดี ประกอบกับการใช้จำนวนโดเมนข่าวที่เผยแพร่ข่าวปลอมและโดเมนข่าวที่เผยแพร่ข่าวจริง ก็เป็นอีก 2 ปัจจัย ที่ช่วยให้การจำแนกข่าวได้ดีขึ้น ข้อมูลที่ได้ทำให้เกิดการจำแนกได้เป็นอย่างดีเห็นได้ชัดจากการพล็อตค่าสหสัมพันธ์และการวิเคราะห์ความสัมพันธ์ระหว่างคุณลักษณะกับประเภทข่าว ทั้ง 5 ลักษณะเด่นที่ผู้วิจัยใช้ส่งผลให้การเรียนรู้ของเรียนรู้ข้อมูลได้อย่างมีประสิทธิภาพดังจะเห็นจากหลายแบบจำลองมีความถูกต้องสูงกว่า 90% สำหรับงานวิจัยนี้เลือก LSTM สำหรับเป็นตัวจำแนกข้อมูลในแอปพลิเคชันเนื่องจากโมเดลอยู่ในกลุ่มได้ผลทดลองด้านประสิทธิภาพสูงสุดและมีความเสถียรสูงสุดจากการวัดประสิทธิภาพแบบไขว้ 10 ชุดข้อมูล



รูปที่ 4 ระบบตรวจข่าวปลอมด้วยการเรียนรู้ของเครื่อง <https://thaidimachine.org>

ผลการพัฒนาระบบตรวจสอบข่าวปลอม ผู้วิจัยได้ทำการพัฒนาและติดตั้งระบบบนคลาวด์ที่ประชาชนสามารถเข้าใช้งานได้ที่ <https://thaidimachine.org> ผลการประเมินระบบโดยผู้ใช้งานจำนวน 34 คน แบ่งการประเมินออกเป็น 3 ด้าน ได้แก่ 1. ด้านการทำงานได้ตามฟังก์ชันของระบบ (Functional Test) 2. ด้านความง่ายต่อการใช้งาน (Usability Test) และ 3. ด้านความปลอดภัยของระบบ (Security Test) เว็บไซต์ตรวจสอบข่าวปลอมสามารถใช้งานได้ง่าย เมื่อประชาชนต้องการตรวจสอบข่าว ทำได้โดยการพิมพ์ข้อความหรือทวิตข้อความลงในช่องข้อมูลแล้วกดตรวจสอบข่าว ระบบจะทำการประมวลผลแล้วแสดงผลการวิเคราะห์ข่าว โดยแสดงผลเป็นผลลัพธ์เป็น “ข่าวจริง” “ข่าวปลอม” หรือ “ข่าวน่าสงสัย” พร้อมมีลิงค์ข่าวที่เกี่ยวข้องให้ผู้ตรวจข่าวมีข้อมูลเพิ่มเติมเกี่ยวกับข่าว นอกจากการวิเคราะห์ข่าว ระบบทำการรวบรวมข่าวปลอม ข่าวที่กำลังมีผู้สนใจ และข่าวภาคีรัฐ ระบบเก็บรวบรวมลิงค์จากเว็บไซต์ต่าง ๆ ไว้โดยอัตโนมัติ เพื่อความสะดวกกับผู้ใช้จะได้ตรวจสอบผ่านเว็บที่น่าเชื่อถือ ด้านความปลอดภัยของเว็บแอปพลิเคชัน ระบบมีฟังก์ชันการล็อกอิน/ลงทะเบียน สำหรับผู้ใช้ที่ประสงค์จะป้อนกลับข้อมูลข่าว ต้องมีการลงทะเบียนไว้ด้วยเพื่อระบุตัวตน ให้มีความมั่นใจในการให้ข้อคิดเห็น และระบบมีการป้องกันการบุกรุกโดยใช้ซอฟต์แวร์แบบเปิด ได้แก่ Fail2ban ป้องกันเซิร์ฟเวอร์ มีไฟร์วอลล์แบบ Uncomplicated Firewall (ufw) กำหนดกฎให้ป้องกัน Hackers และ Bad Bots ตัวอย่างหน้าจอแรกของเว็บตรวจข่าวปลอมแสดงดังรูปที่ 4

หลังจากเปิดการใช้งานผู้วิจัยติดตามผลการใช้งาน พบว่ามีการสืบค้นเฉลี่ยมากกว่าวันละ 1,000 ข้อความที่ได้รับการสืบค้น จากการวิเคราะห์พบว่าจากการสุ่มถามผู้ใช้งานระบบตรวจข่าวปลอม จำนวน 34 คน มีความคิดเห็นเกี่ยวกับในระบบที่ระดับความคิดเห็นในทุกด้านทุกรายการประเมินในระดับ ดีมาก ด้วยค่าเฉลี่ยทุกด้านทุกรายการประเมิน มีค่าเท่ากับ 4.53 และส่วนเบี่ยงเบนมาตรฐาน 0.63 ผลประเมินด้านความปลอดภัยของระบบอยู่ในระดับต่ำสุด มีค่าเฉลี่ยเท่ากับ 4.46 ค่าเบี่ยงเบนมาตรฐานเท่ากับ 0.70 อยู่ในระดับดี เมื่อพิจารณาผลการประเมินด้านการทำงานได้ตามฟังก์ชันของระบบ มีค่าเฉลี่ย 4.52 และส่วนเบี่ยงเบนมาตรฐานเท่ากับ 0.54 อยู่ในระดับดีมาก ผลการประเมินด้านความง่ายต่อการใช้งาน มีค่าเฉลี่ยเท่ากับ 4.59 ส่วนเบี่ยงเบนมาตรฐาน 0.55 อยู่ในระดับดีมาก ผู้ประเมินซึ่งส่วนใหญ่จบการศึกษาสาขาที่เกี่ยวข้องกับคอมพิวเตอร์และเทคโนโลยีสารสนเทศ และมีประสบการณ์การทำงานด้านคอมพิวเตอร์และเทคโนโลยีสารสนเทศโดยตรง ให้ความคิดเห็นในภาพรวมว่าระบบที่ผู้วิจัยพัฒนาขึ้นนี้มีประสิทธิภาพดีมาก มีการจัดรูปแบบได้สวยงาม มีการใช้ฟอนต์ขนาดพอดี และเป็นแอปพลิเคชันที่มีประโยชน์และสามารถใช้งานได้จริง

สรุป

จากปัญหาข่าวปลอมที่มีจำนวนมากในสื่อออนไลน์ ทำให้มีผู้หลงเชื่อข่าวปลอมว่าเป็นข่าวจริง และนำไปแชร์ต่อทำให้เกิดปัญหาอื่นๆ ตามมา ปัจจุบันยังไม่มีเครื่องมือที่จะช่วยตรวจสอบข่าวปลอมภาษาไทย ที่จะช่วยชะลอการแพร่กระจายข่าว การตรวจจับข่าวปลอมเป็นงานที่ยากเนื่องจากข่าวมีความเคลื่อนไหวอย่างมาก งานวิจัยนี้เสนอวิธีการใหม่ที่มีประสิทธิภาพในการจัดการกับข่าวปลอมหรือข้อมูลที่ผิด ผู้วิจัยจึงได้นำเสนอระบบตรวจสอบข่าว

ปลอมภาษาไทยที่ทำงานบนพื้นฐานเทคนิคทางปัญญาประดิษฐ์ โดยเครื่องประมวลผลตรวจสอบข่าวบนคลาวด์และส่งผลการวิเคราะห์ข่าวให้ผู้ใช้งานทราบ

ผู้วิจัยใช้เว็บเบราว์เซอร์รวบรวมข้อมูลสำหรับตัวอย่าง 41,448 ตัวอย่าง และทำการจัดประเภทไว้ล่วงหน้าเป็นข่าวจริง ข่าวปลอม และข่าวน่าสงสัย จำนวนตัวอย่างข้อมูลในแต่ละกลุ่มมีความสมดุลโดยสุ่มเพิ่มเป็น 53,220 ตัวอย่าง ผู้วิจัยใช้วิธีการสอนแบบไขว้ทดสอบจำนวน 10 ชุดข้อมูล (10-Fold Cross Validation) แบบจำลองการเรียนรู้ของเครื่องซึ่งเป็นศาสตร์ด้านปัญญาประดิษฐ์ที่ใช้ในการศึกษา ได้แก่ การถดถอยโลจิสติกส์ (LR) เพื่อนบ้านใกล้เคียง (KNN) นาอิวเบย์ (NB) โครงข่ายประสาทเทียมเพอร์เซ็ปตรอนหลายชั้น (MLP) ซัพพอร์ตเวกเตอร์ (SVM) ป่าสุ่ม (RF) และโครงข่ายประสาทเทียมชนิดหน่วยความจำระยะสั้นแบบยาว (LSTM) ผลการเปรียบเทียบประสิทธิภาพโดยเฉลี่ยของการเรียนรู้ของเครื่องชนิดต่าง ๆ กลุ่มโมเดลที่พบว่าดีที่สุดในงานวิจัยนี้ได้แก่ DT, RF, SVM, LSTM และ MLP ซึ่งเหมาะสำหรับการเรียนรู้ของเครื่องที่จะนำไปใช้เป็นระบบตรวจข่าวปลอม

ระบบตรวจสอบข่าวปลอมบนพื้นฐานด้านปัญญาประดิษฐ์มีการทำงานประกอบด้วย 3 องค์ประกอบหลัก ได้แก่ ส่วนการสืบค้นข้อมูลข่าวบนเว็บ ส่วนการประมวลผลภาษาธรรมชาติ และส่วนการเรียนรู้ของเครื่อง ส่วนแรกที่เป็นการสืบค้นข่าวบนเว็บทำหน้าที่สืบค้นข่าวที่ตรงกับข่าวสืบค้นซึ่งป้อนเข้าระบบโดยผู้ใช้งานที่ประสงค์จะตรวจข่าว ส่วนที่สองจะรับข้อมูลข่าวที่สืบค้นได้จากส่วนแรกและนำข้อมูลข่าวไปประมวลผล โดยทำการตัดคำและสกัดคำที่เป็นลักษณะเด่นของข่าวจริงข่าวปลอม และส่วนที่สามเป็นการเรียนรู้ของเครื่องทำหน้าที่จำแนกข้อมูลที่ได้รับจากส่วนที่สองและแสดงผลเป็น 3 ประเภท ได้แก่ ข่าวจริง ข่าวปลอม และข่าวน่าสงสัย

ระบบจะตรวจข่าวจากการสืบค้นข่าวเผยแพร่ในอินเทอร์เน็ตมาสกัดคำฟิเจอร์หรือลักษณะเด่นของข่าวก่อนส่งไปตัดสินใจ หากข้อความข่าวไม่มีการแชร์หรือเผยแพร่มากนักทำให้ผลการสกัดคำฟิเจอร์ไปตกอยู่ที่กลุ่มน่าสงสัย แต่หากข่าวมีการเผยแพร่จำนวนมากและมีกลุ่มคำที่มักปรากฏในข่าวปลอม เมื่อสกัดคำฟิเจอร์ส่งผลการตัดสินใจได้เป็นข่าวปลอม และมีการแชร์และเผยแพร่เป็นจำนวนมากและมีกลุ่มคำที่มักปรากฏในข่าวจริง เมื่อสกัดคำฟิเจอร์ส่งผลการตัดสินใจว่าเป็นข่าวจริง

ข้อจำกัดของระบบคือระบบยังไม่เข้าใจเนื้อหาข่าวทั้งหมดแบบกว้างอย่างที่มีมนุษย์เข้าใจ ในการวิจัยต่อไปอนาคต อาจจะต้องประยุกต์ใช้การเรียนรู้เชิงลึกประเภท BERT (Bidirectional Encoder Representations from Transformers) เพื่อให้มีความสามารถในการเข้าใจภาษาธรรมชาติ (Natural Language Understanding) ทำอย่างไรให้เครื่องเข้าใจข่าวมากขึ้นเหมือนมนุษย์ เครื่องสามารถทำนายประเภทของข่าวและอธิบายว่าทำไมจึงให้คำตอบหรือการตอบสนองดังกล่าว นอกจากนี้ หากเนื้อหาข่าวประกอบด้วยข้อความ เสียง และวิดีโอ เครื่องต้องวิเคราะห์และตอบสนองอย่างถูกต้องและมีการเสนอข่าวใกล้เคียงความสอดคล้องกันกับข่าวสืบค้นด้วยซ้ำ

กิตติกรรมประกาศ

งานวิจัยนี้อยู่ภายใต้ความร่วมมือระหว่างหน่วยปฏิบัติการวิจัยและพัฒนาทรัพยากรมนุษย์ด้านความรู้ทางดิจิทัลและการรู้เท่าทันสื่อ (DIRU) คณะนิเทศศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย และคณะเทคโนโลยีสารสนเทศและนวัตกรรมดิจิทัล มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ (KMUTNB) บทความนี้เป็นส่วนหนึ่งของโครงการวิจัยเรื่อง “การพัฒนาการตรวจจับข่าวปลอมโดยการเรียนรู้ของเครื่องและการตรวจสอบข้อเท็จจริงของประชาชน” ได้รับการสนับสนุนเงินทุนวิจัยจากกองทุนพัฒนาสื่อปลอดภัยและสร้างสรรค์ ประจำปีงบประมาณ 2562 ได้รับรหัสโครงการ: 2562-T3/0064

เอกสารอ้างอิง

- Ayoub, J., Yang, X. J., & Zhou, F. (2021). Combat COVID-19 infodemic using explainable natural language processing models. *Information Processing and Management*, 58(4). <https://doi.org/10.1016/j.ipm.2021.102569>
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297. <https://doi.org/10.1007/BF00994018>
- Gori, N. (2018). *Machine Learning: A constraint-based approach*. Morgan Kaufmann.
- Guo, G., Wang, H., Bell, D., Bi, Y., & Greer, K. (2003). KNN Model-Based Approach in Classification. In R. Meersman, R., Tari, Z., & Schmidt, D. C. (Eds.), *On The Move to Meaningful Internet Systems 2003: CoopIS, DOA, and ODBASE* (Vol. 2888, pp. 986–996). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-540-39964-3_62
- Hagan, M. T., Demuth, H. B., & Beale, M. H. (2014). *Neural network design*. Martin Hagan.
- Hochreiter, S., & Schmidhuber, J. (1997). Long Short-Term Memory. *Neural Computation*, 9(8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- Kleinbaum, D. G., & Klein, M. (2010). *Logistic Regression: A Self-Learning Text* (3rd ed.). Springer.
- Krallinger, M., Rabal, O., Lourenço, A., Oyarzabal, J., & Valencia, A. (2017). Information Retrieval and Text Mining Technologies for Chemistry. *Chemical Reviews*, 117(12), 7673–7761. <https://doi.org/10.1021/acs.chemrev.6b00851>
- Myles, A. J., Feudale, R. N., Liu, Y., Woody, N. A., & Brown, S. D. (2004). An introduction to decision tree modeling. *Journal of Chemometrics*, 18(6), 275–285. <https://doi.org/10.1002/cem.873>

- Phatthiyaphaibun, W., Chaovavanich, K., Polpanumas, C., Suriyawongkul, A., Lowphansirikul, L., Chormai, P., smeeklai, Tanruangporn, P. P., Charoenchainetr, P., Udomcharoenchaikit, C., Supaseth, Janthong, A., Chaovavanich, K., Korkeat W., Jitchiranant, N., Ninyawee, N., boomsquared, Dubois, Y., nyamakawa, Somsontorn, C., Cody, Prakalphakul, K., Yuankrathok, P., Badger, C., fossabot, hopedataannotations, & pontakornth. (2016). *PyThaiNLP: Thai Natural Language Processing in Python*. <https://doi.org/10.5281/zenodo.4662045>
- Svetnik, V., Liaw, A., Tong, C., Culberson, J. C., Sheridan, R. P., & Feuston, B. P. (2003). Random Forest: A Classification and Regression Tool for Compound Classification and QSAR Modeling. *Journal of Chemical Information and Computer Sciences*, 43(6), 1947–1958. <https://doi.org/10.1021/ci034160g>
- Yager, R. R. (2006). An extension of the naive Bayesian classifier. *Information Sciences*, 176(5), 577–588. <https://doi.org/10.1016/j.ins.2004.12.006>