

การจำแนกข้อความ SMS ภาษาไทย จากมิจฉาชีพ ด้วย Machine Learning Classifying Thai SMS Messages from Scammers using Machine Learning

พีรพล ศรีบุญ¹, พรมนัส วรรณศรี¹, จีรวัด ไร่เจริญ¹, วาสนา ดวงเหมือน¹ และ สุรเทพ แป้นเกิด^{1*}
Peerapon Sribun¹, Pornmanat Wannasri¹, Jerawat Raicharoen¹, Wassana Duangmeun¹ and
Suratep Pangerd^{1*}

¹ สาขาวิชาเทคโนโลยีสารสนเทศและธุรกิจดิจิทัล คณะบริหารธุรกิจ มหาวิทยาลัยเทคโนโลยีราชมงคลกรุงเทพ เลขที่ 2 ถนนนางลิ้นจี่
แขวงทุ่งมหาเมฆ เขตสาทร กรุงเทพมหานคร 10120

¹ Division of Information Technology and Digital Business, Faculty of Business Administration, Rajamangala University of
Technology Bangkok, 2 Nang Linchi Road, Thung Maha Mek, Sathorn, Bangkok 10120, Thailand.

*Corresponding Author: Suratep.p@mail.rmutk.ac.th

Received 27 มีนาคม 2569; Revised 20 เมษายน 2569; Accepted 22 เมษายน 2569

บทคัดย่อ

การวิจัยครั้งนี้มีวัตถุประสงค์ 1) เพื่อการจำแนกข้อความ SMS ภาษาไทยจากมิจฉาชีพด้วย Machine Learning 2) ศึกษาการใช้อัลกอริทึม PyThaiNLP ร่วมกับเทคนิค TF-IDF ในการสกัดคุณลักษณะข้อความ และประยุกต์ใช้เทคนิค SMOTE เพื่อแก้ไขปัญหาความไม่สมดุลของข้อมูล ก่อนนำไปฝึกด้วยอัลกอริทึม Naïve Bayes และ 3) ศึกษาผลของแบบจำลองในการจำแนกข้อความ SMS ภาษาไทยออกเป็น 3 กลุ่ม ได้แก่ ข้อความปกติ ข้อความมิจฉาชีพ และข้อความส่งเสริมการขาย โดยใช้ข้อมูลจำนวน 5,000 รายการ ผลการทดลองด้วยอัลกอริทึม Naïve Bayes พบว่า กลุ่มข้อความปกติมีค่า Precision (ความแม่นยำ) 0.94 Recall (ความระลึก) 0.84 และ F1-Score (ค่าความถ่วงดุล) 0.88 กลุ่มข้อความส่งเสริมการขายมีค่า Precision 0.57 Recall 0.81 และ F1-Score 0.67 และกลุ่มข้อความมิจฉาชีพมีค่า Precision 0.84 Recall 0.91 และ F1-Score 0.87 ซึ่งแสดงให้เห็นถึงประสิทธิภาพที่ดีของแบบจำลอง โดยเฉพาะในการตรวจจับข้อความมิจฉาชีพ นอกจากนี้ การประยุกต์ใช้เทคนิค SMOTE ช่วยเพิ่มประสิทธิภาพของแบบจำลอง โดยให้ค่าความถูกต้อง (Accuracy) สูงสุดร้อยละ 85.00 พร้อมทั้งให้ค่าความแม่นยำ ความระลึก และค่าความถ่วงดุลอยู่ในระดับดีเยี่ยม สรุปได้ว่าแนวทางดังกล่าวสามารถนำไปใช้ในการตรวจจับข้อความ SMS มิจฉาชีพได้อย่างมีประสิทธิภาพ

คำหลัก: การประมวลผลภาษาธรรมชาติ; การจำแนกข้อความ; ข้อความสแปม; Naïve Bayes; การเรียนรู้ของเครื่อง

Abstract

This research aims to: 1) classify Thai SMS messages related to scams using Machine Learning; 2) investigate the use of PyThaiNLP combined with TF-IDF for text feature extraction and apply the SMOTE technique to address class imbalance before training with the Naïve Bayes algorithm; and 3) evaluate the performance of the model in classifying Thai SMS messages into three categories: normal, scam, and promotional messages, using a dataset of 5,000 samples. The experimental results using the Naïve Bayes algorithm show that the normal class achieved a Precision (accuracy of positive predictions) of 0.94, Recall (ability to detect actual positives) of 0.84, and F1-Score (balance between Precision and Recall) of 0.88. The

promotional class obtained a Precision of 0.57, Recall of 0.81, and F1-Score of 0.67, while the scam class achieved a Precision of 0.84, Recall of 0.91, and F1-Score of 0.87. These results indicate that the model performs well, particularly in detecting scam messages. Furthermore, applying the SMOTE technique improved the model performance, achieving the highest Accuracy of 85.00%, with Precision, Recall, and F1-Score at a good level. In conclusion, this approach is effective for Thai SMS scam detection and can be applied in real-world scenarios.

Keywords: Natural Language Processing; Text Classification; Spam SMS; Naïve Bayes; Machine Learning

1. บทนำ

การสื่อสารผ่านโทรศัพท์เคลื่อนที่และบริการส่งข้อความสั้น (Short Message Service: SMS) ถือเป็นช่องทางการสื่อสารที่สำคัญและเข้าถึงประชาชนได้อย่างรวดเร็วติดต่อสื่อสาร หรือการรับส่งข้อมูลใน ระยะทางไกลมีความสะดวกและรวดเร็วมากยิ่งขึ้น ด้วยรูปแบบที่ง่ายต่อการใช้งานผ่านอุปกรณ์ อิเล็กทรอนิกส์ จึงส่งผลให้อัตราการใช้งานเพิ่มสูงขึ้นอย่างรวดเร็ว โดยผู้ใช้งานสามารถพูดคุย แบ่งปัน แลกเปลี่ยน หรือสิ่งที่ตัวเองสนใจได้อย่างไม่มีขีดจำกัดโดยใช้หลักการการประมวลผลภาษาธรรมชาติ(Natural language processing - NLP) และการเรียนรู้ของเครื่อง (Machine learning) ซึ่งเป็นสาขาหนึ่งของปัญญาประดิษฐ์ (Artificial intelligence) เพื่อสร้างขั้นตอนวิธีเพื่อจำแนกประเภทข้อความ ออกจากข้อความทั่วไป การจำแนกประเภทข้อความเป็นหนึ่งในวิธีที่พบได้บ่อยที่สุดในการจัดเรียงข้อความออกเป็นกลุ่มอย่างเป็นระบบ เรียกอีกอย่างว่าการติดแท็กข้อความหรือการจัดหมวดหมู่ โปรแกรมจำแนกประเภทข้อความสามารถวิเคราะห์ข้อความได้โดยใช้การประมวลผลภาษาธรรมชาติ (NLP)

ในยุคดิจิทัลปัจจุบัน โทรศัพท์เคลื่อนที่และบริการส่งข้อความสั้น (Short Message Service: SMS) ถือเป็นช่องทางการสื่อสารที่สำคัญและเข้าถึงประชาชนได้อย่างรวดเร็ว อย่างไรก็ตาม ความสะดวกรวดเร็วนี้ได้ถูกนำมาใช้เป็นเครื่องมือโดยกลุ่มมิจฉาชีพและผู้ไม่ประสงค์ดีในการหลอกลวงประชาชน (Scam) เช่น การส่งลิงก์ปลอมเพื่อขโมยข้อมูลส่วนบุคคล (Phishing) ข้อความ SMS (เบอร์โทรศัพท์มือถือ): มักจะส่งมาจากเบอร์แปลก หรือบางครั้งปลอมเป็นชื่อหน่วยงานที่ท่านคุ้นเคย เช่น ธนาคาร, บริษัทขนส่ง, หรือแม้แต่การไฟฟ้า/ประปา [1]

การหลอกลวงให้โอนเงิน การโฆษณาเว็บพนันออนไลน์ รวมถึงการส่งข้อความส่งเสริมการขาย (Promotional Message) ที่สร้างความรำคาญให้แก่ผู้รับ ปัญหาดังกล่าวส่งผลกระทบต่อความมั่นคงปลอดภัยทางไซเบอร์และก่อให้เกิดความสูญเสียทางทรัพย์สินจำนวนมากในประเทศไทย แม้ว่าผู้ให้บริการเครือข่ายจะมีการบล็อกหมายเลขโทรศัพท์ที่น่าสงสัย แต่กลุ่มมิจฉาชีพมักจะปรับเปลี่ยนรูปแบบข้อความและหมายเลขผู้ส่งอยู่เสมอ ทำให้การรับมือด้วยวิธีดั้งเดิมไม่เพียงพอต่อการป้องกัน

2. วิธีดำเนินงานวิจัย

การวิจัยครั้งนี้เป็นการวิจัยเชิงทดลอง (Experimental Research) โดยมีวัตถุประสงค์เพื่อพัฒนาแบบจำลองการเรียนรู้ของเครื่อง (Machine Learning) สำหรับจำแนกประเภทข้อความ SMS ที่เป็นภาษาไทย ซึ่งผู้วิจัยได้กำหนดขั้นตอนการดำเนินงานออกเป็น 5 ขั้นตอนหลัก ดังนี้โดยใช้ อัลกอริทึม PyThaiNLP ร่วมกับเทคนิค TF-IDF เพื่อตัดคำและแปลงตัวหนังสือให้เป็นตัวเลข เพื่อให้คอมพิวเตอร์คำนวณได้และประยุกต์ใช้เทคนิค SMOTE เพื่อ แก้ไขปัญหาความไม่สมดุลของข้อมูล ดังรูปที่ 1 และมีรายละเอียดการดำเนินการวิจัยดังนี้

2.1 การจัดการข้อมูลและการเตรียมความพร้อมล่วงหน้า (Data Management and Preprocessing)

ชุดข้อมูลที่ใช้ในการวิจัยประกอบด้วยข้อความสั้น (SMS) ภาษาไทยจำนวน 5,000 รายการ โดยมีการกำหนดป้ายกำกับ (Label) เพื่อแบ่งประเภทของข้อความออกเป็น 3 คลาส (Class) ได้แก่ ข้อความปกติ (Normal) จำนวน 3,217 รายการ ข้อความมิจฉาชีพ (Scam) จำนวน 1,184 รายการ และข้อความส่งเสริมการขาย (Promo) จำนวน 599

รายการ จากสถิติเบื้องต้นพบว่าชุดข้อมูลมีลักษณะไม่สมดุล (Imbalanced Data) อย่างชัดเจน ซึ่งสอดคล้องกับ พฤติกรรมการรับข้อความในสถานการณ์จริง

2.2. การทำความสะอาดและประมวลผลข้อความเบื้องต้น (Text Preprocessing)

คอมพิวเตอร์ไม่สามารถประมวลผลข้อความภาษาไทย ในรูปแบบประโยคยาวได้โดยตรง ผู้วิจัยจึงนำคลังข้อมูลและไลบรารี PyThaiNLP มาใช้ในกระบวนการเตรียมข้อมูล โดยมีขั้นตอนดังนี้ 1) การแปลงตัวอักษรภาษาอังกฤษทั้งหมดให้เป็นตัวพิมพ์เล็ก (Case Folding) 2) การตัดคำ (Tokenization) ด้วยเครื่องมือเอนจิน (Engine) ที่ชื่อว่า 'newmm' ซึ่งอาศัยหลักการจับคู่คำที่ยาวที่สุดตามพจนานุกรม 3) การกรองคำฟุ่มเฟือย (Stopwords Removal) โดยตัดคำที่ไม่มีผลต่อความหมายหลักของประโยคออก และ 4) การทำความสะอาดด้วยนิพจน์พหุคูณ (Regular Expression) เพื่อเก็บไว้เฉพาะตัวอักษรภาษาไทย ภาษาอังกฤษ และตัวเลข โดยตัดเครื่องหมายวรรคตอนและสัญลักษณ์พิเศษทิ้งไป

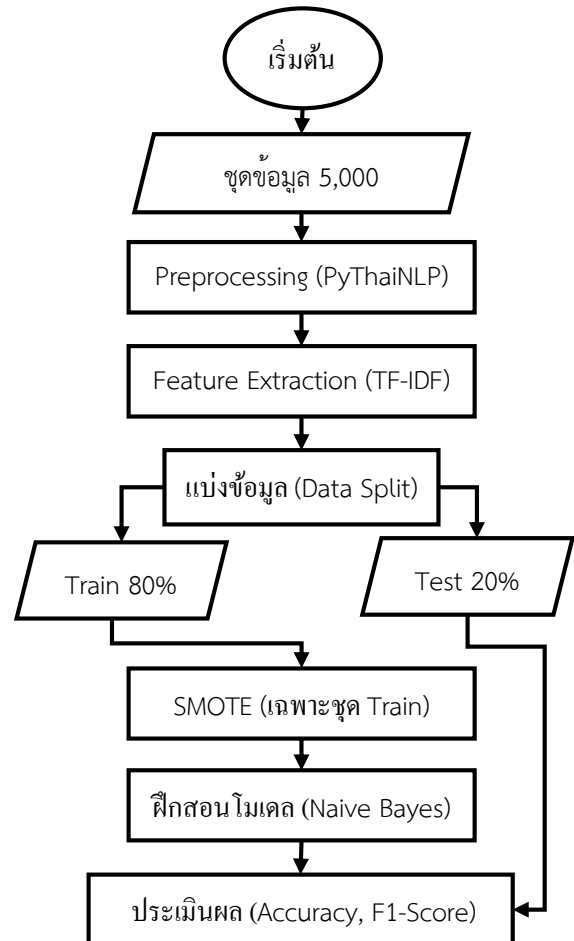
2.3. การสกัดคุณลักษณะข้อความและการแบ่งข้อมูล (Feature Extraction and Data Splitting)

หลังจากการทำความสะอาดข้อความ ข้อมูลจะถูกนำมาแปลงให้อยู่ในรูปแบบเวกเตอร์ตัวเลข (Numerical Vector) ด้วยเทคนิค TF-IDF (Term Frequency-Inverse Document Frequency) เพื่อคำนวณค่าน้ำหนักความสำคัญของแต่ละคำศัพท์ในชุดข้อมูล โดยคำที่ปรากฏบ่อยในข้อความประเภทใดประเภทหนึ่งแต่ไม่ค่อยพบในข้อความประเภทอื่น จะได้รับค่าน้ำหนักที่สูง จากนั้นผู้วิจัยได้ทำการแบ่งชุดข้อมูล (Train-Test Split) ออกเป็น 2 ส่วน คือ ชุดข้อมูลสำหรับฝึกสอน (Training Set) ร้อยละ 80 และชุดข้อมูลสำหรับทดสอบ (Testing Set) ร้อยละ 20

2.4. การจัดการชุดข้อมูลที่ไม่สมดุล (Data Augmentation with SMOTE)

เนื่องจากชุดข้อมูลสำหรับฝึกสอนมีสัดส่วนของข้อความปกติมากกว่าข้อความมีฉ้อฉลและข้อความส่งเสริมการขายอย่างมาก ซึ่งอาจส่งผลให้แบบจำลองเกิดการเรียนรู้ที่ลำเอียง (Bias) ผู้วิจัยจึงประยุกต์ใช้เทคนิค SMOTE (Synthetic Minority Over-sampling Technique) กับชุด

ข้อมูลสำหรับฝึกสอน (Training Set) เท่านั้น เพื่อสร้างข้อมูลสังเคราะห์สำหรับคลาสที่เป็นชนกลุ่มน้อย (Minority Class) ให้มีจำนวนเทียบเท่ากับคลาสหลัก ทำให้แบบจำลองสามารถเรียนรู้คุณลักษณะของข้อความทุกประเภทได้อย่างเท่าเทียมและมีประสิทธิภาพมากขึ้น



รูปที่ 1 แผนภาพสรุปขั้นตอนการดำเนินงานวิจัย

2.5 การสร้างแบบจำลองและการประเมินผล (Model Training and Evaluation)

งานวิจัยนี้เลือกใช้อัลกอริทึมการเรียนรู้ของเครื่องแบบ โนอีฟเบย์ (Multinomial Naïve Bayes) ในการสร้างแบบจำลองเพื่อจำแนกประเภทข้อความ เนื่องจากเป็นอัลกอริทึมที่เหมาะสมกับข้อมูลประเภทข้อความและอาศัยหลักการทางสถิติความน่าจะเป็นในการจัดหมวดหมู่ หลังจากการฝึกสอนเสร็จสิ้น แบบจำลองจะถูกนำไปประเมินประสิทธิภาพกับชุดข้อมูลสำหรับทดสอบ (Testing Set) โดยใช้ตัวชี้วัดมาตรฐาน ได้แก่ ค่าความแม่นยำรวม (Accuracy) ค่าความแม่นยำ (Precision) ค่าความถูกต้อง (Recall) และค่าเอฟวัน (F1-Score) รวมถึงการวิเคราะห์

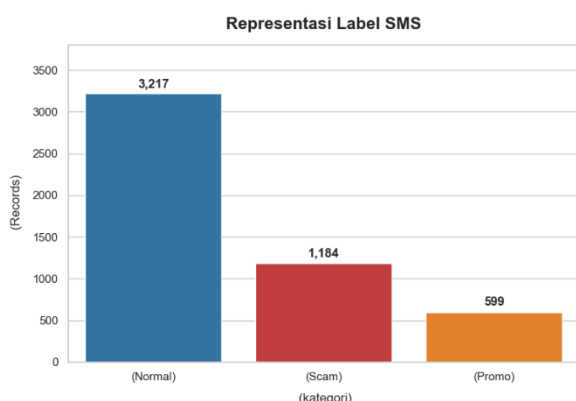
ค่าเฉลี่ยแบบมหภาค (Macro Average) และแบบถ่วงน้ำหนัก (Weighted Average) ควบคู่ไปกับการใช้เมทริกซ์ความสับสน (Confusion Matrix) เพื่อตรวจสอบการทำนายที่ผิดพลาดในแต่ละคลาสอย่างละเอียด (Workflow)

3. ผลและการอภิปรายผลการวิจัย

การวิจัยนี้ได้นำชุดข้อมูลข้อความสั้น (SMS) ภาษาไทย จำนวน 5,000 รายการ มาผ่านกระบวนการทำความสะอาดข้อความและสกัดคุณลักษณะด้วยเทคนิค TF-IDF ก่อนนำไปปรับสมดุลข้อมูลด้วยเทคนิค SMOTE และฝึกสอนด้วยอัลกอริทึม Multinomial Naïve Bayes โดยผลการดำเนินงานวิจัยสามารถแบ่งออกได้ดังนี้

3.1 ข้อมูลทั่วไปของชุดข้อมูลก่อนการปรับสมดุล

จากการสำรวจชุดข้อมูลข้อความตั้งต้น พบว่ามีการกระจายตัวของคลาสที่ไม่สมดุลกัน (Imbalanced Data) อย่างเห็นได้ชัด โดยประกอบด้วย ข้อความปกติ (Normal) จำนวน 3,217 รายการ คิดเป็นร้อยละ 64.34 ของข้อมูลทั้งหมด ตามด้วยข้อความมิจฉาชีพ (Scam) จำนวน 1,184 รายการ (ร้อยละ 23.68) และข้อความส่งเสริมการขาย (Promo) จำนวน 599 รายการ (ร้อยละ 11.98) ดังแสดงในรูปที่ 1 ซึ่งความไม่สมดุลนี้เป็นเหตุผลสำคัญที่ผู้วิจัยต้องนำเทคนิค SMOTE มาประยุกต์ใช้เพื่อป้องกันการเกิดความลำเอียงของแบบจำลอง ดังรูปที่ 2



รูปที่ 2 จำนวนรายการข้อความ SMS แบ่งตามประเภทในชุดข้อมูลตั้งต้น

3.2 ประสิทธิภาพของแบบจำลองในการจำแนก

ข้อความ

เมื่อนำชุดข้อมูลที่ผ่านการปรับสมดุลด้วยเทคนิค SMOTE มาฝึกสอนและทดสอบกับแบบจำลอง

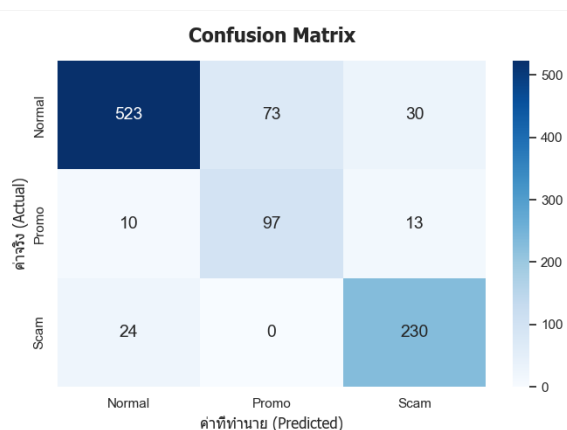
Multinomial Naïve Bayes [3] พบว่าแบบจำลองมีประสิทธิภาพในการจำแนกประเภทข้อความได้เป็นอย่างดี โดยมีค่าความแม่นยำรวม (Accuracy) สูงถึงร้อยละ 85.00 เมื่อพิจารณาค่าทางสถิติของแต่ละคลาส (ตารางที่ 1) พบว่าแบบจำลองสามารถทำนายข้อความปกติและข้อความมิจฉาชีพได้อย่างแม่นยำ โดยมีค่าเอฟวัน (F1-Score) อยู่ที่ 0.88 และ 0.87 ตามลำดับ นอกจากนี้ เมื่อพิจารณาค่าเฉลี่ยแบบมหภาค (Macro Average) และค่าเฉลี่ยแบบถ่วงน้ำหนัก (Weighted Average) ซึ่งมีค่าเท่ากับ 0.81 และ 0.86 ตามลำดับ ดังตารางที่ 1

ตารางที่ 1 สรุปผลประเมินประสิทธิภาพของแบบจำลอง

	Precision	Recall	F1-Score
normal	0.94	0.84	0.88
promo	0.57	0.81	0.67
scam	0.84	0.91	0.87
Macro avg	0.78	0.85	0.81
Weighted avg	0.87	0.85	0.86
Accuracy : 85.00%			

3.3 การวิเคราะห์ความสับสนในการทำนายของแบบจำลอง (Confusion Matrix Analysis)

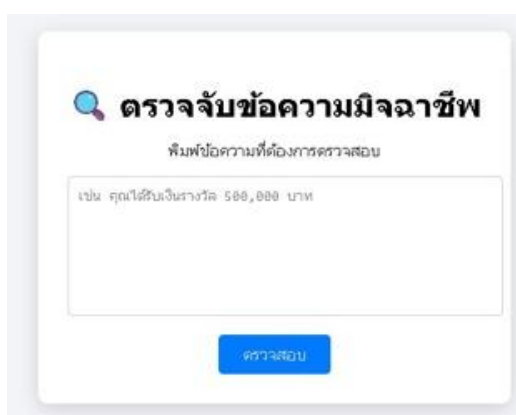
วิเคราะห์ข้อผิดพลาดของแบบจำลองเชิงลึก ผู้วิจัยได้นำผลการทำนายมาแสดงในรูปแบบเมทริกซ์ความสับสน (Confusion Matrix) ดังรูปที่ 3 จากเมทริกซ์พบว่าแบบจำลองสามารถทำนายข้อความแต่ละประเภทได้ถูกต้องเป็นส่วนใหญ่ (พิจารณาจากแนวทแยงมุมหลักของตาราง) อย่างไรก็ตาม ยังพบการทำนายผิดพลาดที่น่าสนใจ กล่าวคือ มีความสับสนระหว่างข้อความปกติและข้อความส่งเสริมการขายอยู่บ้าง จำนวน 10 รายการ ซึ่งสามารถอธิบายได้ว่าข้อความทั้งสองประเภทยังอาจมีการใช้กลุ่มคำศัพท์ (Vocabulary) ที่ใกล้เคียงกัน เช่น คำว่า "ระบบ" "เวลา" หรือ "คุณ" ทำให้แบบจำลองที่พิจารณาความถี่ของคำ (TF-IDF) เกิดความสับสนได้ในบางกรณี ทว่าข้อผิดพลาดดังกล่าวมีจำนวนน้อยมากเมื่อเทียบกับการทำนายที่ถูกต้องทั้งหมด



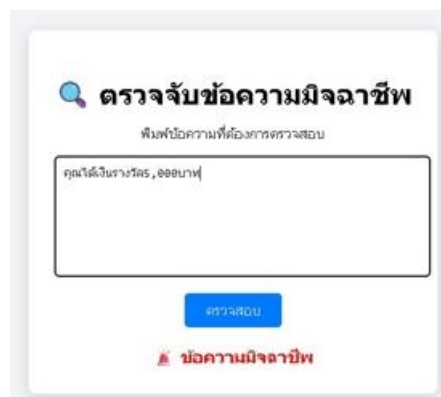
รูปที่ 3 เมทริกซ์ความสับสน (Confusion Matrix) ของแบบจำลอง

3.4 ผลการประเมินประสิทธิภาพ

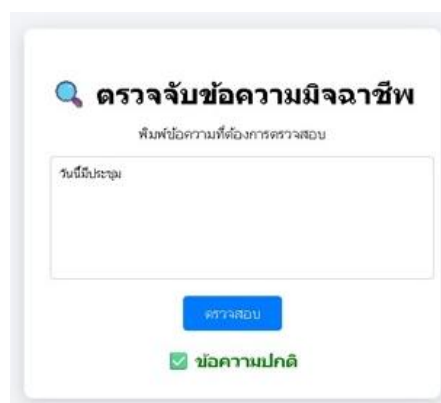
ผลการประเมินจากการทดลองแสดงให้เห็นว่าแบบจำลอง Multinomial Naïve Bayes ที่ผ่านการปรับสมดุลข้อมูลด้วยเทคนิค SMOTE [2] มีประสิทธิภาพในการจำแนกข้อความ SMS ภาษาไทยได้อย่างมีประสิทธิภาพ โดยสามารถจำแนกข้อความออกเป็น 3 ประเภทได้อย่างแม่นยำ ทั้งข้อความปกติ ข้อความมิจฉาชีพ และข้อความส่งเสริมการขาย ทั้งนี้แบบจำลองมีค่าความถูกต้อง (Accuracy) อยู่ที่ร้อยละ 85.00 และมีค่า F1-Score ของแต่ละคลาสอยู่ในระดับสูง สะท้อนถึงความสามารถในการทำนายที่ดี นอกจากนี้ ค่า Macro Average และ Weighted Average ยังอยู่ในเกณฑ์ดี แสดงถึงความสมดุลของประสิทธิภาพในทุกกลุ่มข้อมูล จึงสามารถนำไปประยุกต์ใช้เป็นระบบคัดกรองข้อความเพื่อเพิ่มความปลอดภัยให้กับผู้ใช้งานได้อย่างเหมาะสม ดังรูปที่ 4-6



รูปที่ 4 กล่องใส่ข้อความ



รูปที่ 5 ข้อความมิจฉาชีพ



รูปที่ 6 ข้อความปกติ

3.5 อภิปรายผลการวิจัย

จากผลการทดลองพบว่าเทคนิค SMOTE มีบทบาทสำคัญในการช่วยแก้ไขปัญหาค่าความไม่สมดุลของข้อมูล ส่งผลให้แบบจำลองสามารถเรียนรู้ข้อมูลในกลุ่มส่วนน้อยได้ดีขึ้น ซึ่งสอดคล้องกับงานวิจัยที่ศึกษาการเปรียบเทียบประสิทธิภาพของตัวแบบการเรียนรู้ของเครื่อง [3] และทฤษฎีการจัดการข้อมูลที่ไม่สมดุลของ Chawla et al. [4] ที่ระบุว่า การเพิ่มข้อมูลตัวอย่างสามารถช่วยเพิ่มประสิทธิภาพของแบบจำลองได้อย่างมีนัยสำคัญ นอกจากนี้ยังสอดคล้องกับแนวทางการจำแนกข้อความสนทนา ที่ชี้ให้เห็นว่าการบูรณาการเทคนิคการเพิ่มข้อมูล (Data Augmentation) เข้ากับแบบจำลองการเรียนรู้ สามารถยกระดับความแม่นยำในการคัดกรองข้อความได้อย่างมีประสิทธิภาพ [5] ในส่วนของการใช้ TF-IDF ร่วมกับเครื่องมือประมวลผลภาษาธรรมชาติ PyThaiNLP พบว่าสามารถช่วยให้แบบจำลองแยกแยะความสำคัญของคำในข้อความได้อย่างมีประสิทธิภาพ อย่างไรก็ตาม ยังพบข้อจำกัดในกรณีที่ข้อความต่างประเภทมีการใช้คำศัพท์ที่ใกล้เคียงกัน เช่น ข้อความปกติและข้อความส่งเสริมการขาย ซึ่งส่งผลให้เกิดความสับสนในการ

จำแนกบางส่วน สอดคล้องกับแนวคิดของ Wu et al. [6] ที่อธิบายว่าการให้น้ำหนักคำอาจไม่สามารถสะท้อนความหมายเชิงบริบทได้ทั้งหมด

นอกจากนี้ การเลือกใช้อัลกอริทึม Naïve Bayes ซึ่งมีข้อดีด้านความเรียบง่ายและประสิทธิภาพในการประมวลผลข้อความ ยังให้ผลลัพธ์ที่ดีและสอดคล้องกับการประยุกต์ใช้เทคนิคเหมืองข้อมูลเพื่อการจำแนกข้อความภาษาไทย [7] อย่างไรก็ตาม แบบจำลองดังกล่าวยังมีข้อจำกัดในการพิจารณาความสัมพันธ์ของคำในเชิงลำดับ ทำให้ไม่สามารถเข้าใจบริบทของประโยคได้อย่างลึกซึ้ง

ยิ่งไปกว่านั้น เมื่อพิจารณาในด้านการนำไปใช้ประโยชน์พบว่า พฤติกรรมการรับส่งข้อความ SMS ยังคงเป็นช่องทางที่เข้าถึงผู้ใช้งานโทรศัพท์เคลื่อนที่ได้โดยตรง [8] ผลลัพธ์จากการพัฒนาแบบจำลองในงานวิจัยนี้ จึงสามารถนำไปต่อยอดเป็นระบบคัดกรองเบื้องต้นเพื่อรับมือกับกลโกงในโลกไซเบอร์ [9] และข้อความหลอกลวงที่แนบลิงก์อันตราย ตลอดจนเป็นเครื่องมือสำคัญในการยกระดับความปลอดภัยและป้องกันประชาชนจากการตกเป็นเหยื่อของกลุ่มอาชญากรทางเทคโนโลยี เช่น แก๊งคอลเซ็นเตอร์ [10] หรือการหลอกลวงให้โอนเงิน (Romance Scam) [11] ได้อย่างเป็นรูปธรรม

4. สรุปผลการวิจัย

งานวิจัยนี้ได้นำเสนอการพัฒนาแบบจำลองปัญญาประดิษฐ์เพื่อจำแนกประเภทข้อความสั้น (SMS) ภาษาไทย โดยแบ่งออกเป็น 3 ประเภท ได้แก่ ข้อความปกติ ข้อความมิจฉาชีพ และข้อความส่งเสริมการขาย จากการทดลองพบว่า การประยุกต์ใช้เครื่องมือ PyThaiNLP ร่วมกับเทคนิคการหาค่าน้ำหนักคำแบบ TF-IDF สามารถสกัดคุณลักษณะของข้อความภาษาไทยได้อย่างมีประสิทธิภาพ นอกจากนี้ การนำเทคนิค SMOTE มาใช้เพื่อแก้ปัญหาความไม่สมดุลของชุดข้อมูล (Imbalanced Data) ถือเป็นปัจจัยสำคัญที่ช่วยลดความลำเอียงของแบบจำลอง ทำให้แบบจำลอง Multinomial Naïve Bayes สามารถเรียนรู้และทำนายคลาสที่เป็นกลุ่มย่อย (ข้อความมิจฉาชีพและข้อความส่งเสริมการขาย) ได้อย่างแม่นยำยิ่งขึ้น ผลการประเมินประสิทธิภาพโดยรวมแสดงให้เห็นว่า แบบจำลองที่นำเสนอมีความแม่นยำสูงและมีอัตราการทำนายผิดพลาดต่ำ จึงสรุปได้ว่ากระบวนการและอัลกอริทึมที่เลือกใช้มีความ

เหมาะสม และสามารถนำไปประยุกต์ใช้เป็นระบบคัดกรองข้อความเบื้องต้นเพื่อแจ้งเตือนภัยแก่ผู้ใช้งานโทรศัพท์เคลื่อนที่ได้จริง ซึ่งจะสามารถแยกแยะข้อความ (SMS) โดยผ่าน Web Application ได้

เอกสารอ้างอิง

- [1] ธ. ไทยพาณิชย์, “คลิกเดียวชีวิตเปลี่ยน! วิธีรับมือ ‘ลิงก์แปลกปลอม’ และ ‘ข้อความหลอกลวง’, ” [ออนไลน์]. <https://www.scb.co.th/th/personal-banking/fraud-fighter/update-fraud/sms-one-click>. (เข้าถึงเมื่อ: 9 เมษายน 2569).
- [2] W. Phatthiyaphaibun et al., "PyThaiNLP: Thai Natural Language Processing in Python," in *Proc. 3rd Workshop for Natural Language Processing Open Source Software (NLP-OSS 2023)*, 2023, pp. 25–36.
- [3] ศวิตา ทองขุนวงศ์ และ ภัคพล สวัสดิ์มงคล, "การเปรียบเทียบประสิทธิภาพของตัวแบบการเรียนรู้ของเครื่องสำหรับการจำแนกผู้ป่วยโรคมะเร็งปอด," *วารสารวิจัย มทร. กรุงเทพฯ*, ปีที่ 18, ฉบับที่ 1, หน้า 33–42, มกราคม-มิถุนายน 2567.
- [4] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *J. Artif. Intell. Res.*, vol. 16, pp. 321–357, Jun. 2002.
- [5] นันทวัฒน์ หล้าชีว, "กรอบพัฒนาวิธีการจำแนกประเภทข้อความสนทนาด้วยเทคนิคการเรียนรู้เชิงลึกร่วมกับเทคนิคการเพิ่มข้อมูล," วิทยานิพนธ์ วท.ม. (วิทยาการคอมพิวเตอร์), มหาวิทยาลัยธรรมศาสตร์, กรุงเทพฯ, 2566.
- [6] H. C. Wu, R. W. P. Luk, K. F. Wong, and K. L. Kwok, "Interpreting TF-IDF term weights as making relevance decisions," *ACM Trans. Inf. Syst.*, vol. 26, no. 3, pp. 1–37, Jun. 2008.
- [7] งามอาจ อุ่นอนันต์ และ พยุง มีสัง, "การจำแนกความน่าเชื่อถือของเว็บไซต์แหล่งข่าวภาษาไทยโดยใช้เทคนิคการทำเหมืองข้อมูล," *วารสารวิชาการมหาวิทยาลัยอีสเทิร์นเอเซีย ฉบับวิทยาศาสตร์และเทคโนโลยี*, ปีที่ 14, ฉบับที่ 2, หน้า 101–116, พฤษภาคม-สิงหาคม 2563.

- [8] วสันต์ เจริญทองตระกูล, "พฤติกรรมการรับส่งข้อความสั้น (SMS) ทางโทรศัพท์เคลื่อนที่ของนักเรียน นิสิต นักศึกษา ในกรุงเทพมหานคร," วิทยานิพนธ์ ว.ม. (สื่อสารมวลชน), มหาวิทยาลัยธรรมศาสตร์, กรุงเทพฯ, 2546.
- [9] สุปราณี วงษ์แสงจันทร์, ประภาพร กุลลัมรัตน์ชัย และ พิมล จงวรรณท์, "การรับมือกับกลโกงในโลกไซเบอร์," *วารสารวิชาการมหาวิทยาลัยอีสเทิร์นเอเซีย ฉบับวิทยาศาสตร์และเทคโนโลยี*, ปีที่ 19, ฉบับที่ 2, หน้า 14-26, พฤษภาคม-สิงหาคม 2568.
- [10] สิริลักข์ เมืองนิล, ศุภกร ปุณญฤทธิ์ และ สุธัญญ์ กัลยะจิตร, "การป้องกันการตกเป็นเหยื่อแก๊งคอลเซ็นเตอร์," *วารสารกระบวนการยุติธรรม*, ปีที่ 18, ฉบับที่ 3, หน้า 1-22, พฤษภาคม-สิงหาคม 2568.
- [11] ชวัลวิทย์ โสภาศิริทรัพย์, สุธัญญ์ กัลยะจิตร และ พรชุตี นาคพงษ์, "แนวทางป้องกันการตกเป็นเหยื่อหลอกรักออนไลน์ (Romance Scam)," *วารสารกระบวนการยุติธรรม*, ปีที่ 17, ฉบับที่ 3, หน้า 1-16, กันยายน-ธันวาคม 2567.