

ระบบประเมินความน่าเชื่อถือข่าวสารจากทวีตเตอร์แบบเรียลไทม์ Real Time Credibility Evaluation System for News on Twitter

นำพล เพชรผดุง^{1*} และกานดา สายแก้ว²

^{1*, 2} ภาควิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ มหาวิทยาลัยขอนแก่น อ.เมือง จ.ขอนแก่น 40002

Email: krunapon@kku.ac.th

บทคัดย่อ

ทวีตเตอร์เป็นสื่อโซเชียลมีเดียหนึ่งที่ได้รับคามนิยมในการใช้งานมาก ผู้ใช้ทวีตเตอร์มักใช้ทวีตเตอร์เพื่อติดตามและแบ่งปันข่าวล่าสุดเพราะทวีตเตอร์สามารถกระจายข่าวให้กับผู้ใช้ได้อย่างรวดเร็ว แต่อย่างไรก็ตาม ข่าวที่กระจายอาจเป็นข่าวลือที่เป็นเท็จ ซึ่งก็จะทำให้ผู้คนจำนวนมากรู้สึกตกใจ บทความนี้จึงได้นำเสนอเครื่องมือที่ใช้ในการประเมินความน่าเชื่อถือของข่าวบนทวีตเตอร์โดยสร้างโมเดลจำแนกข้อมูลของการวิเคราะห์คุณลักษณะของทวีตเตอร์ในการประเมินความน่าเชื่อถือของข่าว ผลการทดลองกับข่าวทวีตภาษาไทยพบว่าเครื่องมือประเมินความน่าเชื่อถือของข่าวบนทวีตเตอร์ที่บทความนี้เสนอและพัฒนา มีความแม่นยำมากกว่าเครื่องมือที่มีอยู่แล้วประมาณร้อยละ 24.2 และใช้เวลาน้อยกว่าประมาณร้อยละ 90.35

คำสำคัญ : สื่อสังคมออนไลน์, ทวีตเตอร์, ความน่าเชื่อถือของข้อมูล

Abstract

Twitter is one of popular social media nowadays. Twitter users usually read posts and share news because it can spread information quickly to users. However, if the distributed news is a false rumor, this may cause a large number of people to feel shocked. This paper thus proposes a credibility evaluation tool for news on Twitter that can achieve accuracy using classification model development based on Twitter features analysis. Our experimental result on Thai tweet news shows that our system outperforms the existing tool with the increased accuracy about 24.2% while requiring processing time less than that of the existing tool about 90.35%

Keywords: Social media, Twitter, Information credibility

1. บทนำ

ทวีตเตอร์เป็นสื่อสังคมออนไลน์ประเภทหนึ่งที่ได้รับคามนิยมมาก บริษัททวีตเตอร์จำกัดได้ให้ข้อมูลในเดือนมีนาคม 2558 ว่ามีผู้ใช้ที่เข้ามาใช้ทวีตเตอร์อย่างน้อยหนึ่งครั้งต่อเดือนมีจำนวนมากถึงประมาณ 300 ล้านคน (ทวีตเตอร์, 2015) ทวีตเตอร์จะอนุญาตให้ผู้ใช้ทวีตเตอร์โพสต์ข้อความสั้นที่มีความยาวไม่เกิน 140 ตัวอักษร ข้อความดังกล่าวเรียกว่าทวีต และเปิดเผยว่ามีการโพสต์ทวีตจำนวน 500 ล้านทวีตต่อวัน (ทวีตเตอร์, 2015) ทวีตเตอร์เป็นสื่ออีกชนิดหนึ่งที่เมื่อมีข่าวหรือเหตุการณ์สำคัญต่าง ๆ จะถูกใช้ในการกระจายไปให้ผู้คนที่ได้รับรู้ได้รวดเร็วและเรียลไทม์ แต่ถ้าหากข่าวที่ถูกกระจายออกไปนั้นเป็นข่าวหลอกลวง เช่น ข่าวเกี่ยวกับการแจ้งเตือนภัยพิบัติที่ไม่ได้เกิดขึ้นจริง

หากถูกส่งต่อ ๆ กันโดยทุกคนคิดว่าเป็นเรื่องจริงย่อมทำให้ผู้คนเกิดความแตกตื่น และเกิดเหตุการณ์วุ่นวายขึ้นได้ จากการค้นคว้าเพื่อหาเครื่องมือที่สามารถใช้ประเมินความน่าเชื่อถือของข่าวจากสื่อสังคมออนไลน์พบว่าเครื่องมือที่ชื่อ TweetCred (กัปตา เอ และคณะ, 2014) เมื่อนำมาใช้ประเมินความน่าเชื่อถือทวีตข่าวที่เป็นภาษาไทยซึ่งเป็นข้อมูลชุดทดสอบจำนวน 500 ทวีต ความน่าเชื่อถือที่ได้พบว่ามีความถูกต้องอยู่ที่ร้อยละ 66.40 ซึ่งน่าจะมาจากข้อจำกัดด้านข้อมูลที่ใช้สอนระบบหรือคุณลักษณะที่เลือกใช้นั้นเหมาะสมกับข่าวสารที่เป็นภาษาอังกฤษเป็นหลัก งานวิจัยนี้จึงเสนอวิธีปรับปรุงความแม่นยำในการประเมินความน่าเชื่อถือของข่าวจากทวีตเตอร์ที่เป็นภาษาไทยโดยเลือกคุณลักษณะที่นำมาทดลองสร้างโมเดลจำแนกข้อมูลสำหรับใช้ประเมินความน่าเชื่อถือของข่าวจากทวีตเตอร์ให้เหมาะสมกับทวีตข่าวภาษาไทย และทดลองสร้างโมเดลจำแนกข้อมูลโดยใช้ 2 อัลกอริทึมได้แก่ ซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine) และ ออเดอร์ลอจิสติกส์ (Ordered Logistic) แล้วเลือกโมเดลจำแนกข้อมูลที่ให้ผลการจำแนกตามความน่าเชื่อถือที่ถูกต้องแม่นยำกว่ามาพัฒนาเป็นเครื่องมือประเมินความน่าเชื่อถือข่าวจากทวีตเตอร์แบบเรียลไทม์สำหรับผู้ใช้งานทวีตเตอร์ ให้สามารถใช้งานได้ในรูปแบบของโปรแกรมโครมเอกซ์เทนชัน (Chrome extension) ซึ่งดาวน์โหลดได้จาก <http://bit.ly/thaitweetcred>

2. วัตถุประสงค์

เพื่อพัฒนาระบบประเมินความน่าเชื่อถือข่าวสารจากทวีตเตอร์แบบเรียลไทม์

3. งานวิจัยที่เกี่ยวข้อง

จากการศึกษางานวิจัยที่เกี่ยวข้องกับการประเมินความน่าเชื่อถือของข่าวจากสื่อสังคมออนไลน์พบว่า (มอริส และคณะ, 2012) ได้ทำแบบสอบถามสำรวจความคิดเห็นของผู้ใช้ทวีตเตอร์เกี่ยวกับความน่าเชื่อถือของทวีต พบว่าผู้ใช้รู้สึกยากมากในการประเมินความน่าเชื่อถือโดยพิจารณาเฉพาะคุณลักษณะของเนื้อหาข่าวเท่านั้น จึงได้แนะนำกลยุทธ์ให้ผู้เขียนทวีตข้อความที่จะทำให้มีความน่าเชื่อถือมากขึ้นและกลยุทธ์ให้เครื่องมือที่ใช้ในการค้นหาทวีตแสดงข้อมูลทวีตเตอร์เพื่อเพิ่มความน่าเชื่อถือ (ดอนโนแวน และคณะ, 2012) ทำการวิเคราะห์คุณลักษณะต่าง ๆ เช่น แฮชแท็ก (Hash tag) หรือ การอ้างถึง (Mention) ในทวีต เป็นต้น ว่ามีค่าน้ำหนักมากน้อยเพียงใด ที่จะช่วยให้การประเมินความน่าเชื่อถือนั้นถูกต้องแม่นยำและมีการกระจายตัวอย่างไรในกลุ่มข้อมูลที่มีความน่าเชื่อถือและกลุ่มข้อมูลที่ไม่น่าเชื่อถือ (ทอมสัน และคณะ, 2012) นำเสนอการวิเคราะห์ความน่าเชื่อถือของแหล่งข่าวต่าง ๆ เช่น สำนักข่าว สถาบัน องค์กรเอกชน และบุคคลโดยแยกตามอาชีพที่เผยแพร่ข่าวที่เกี่ยวข้องกับเหตุการณ์ภัยพิบัติที่เมืองฟูกุชิม่า ประเทศญี่ปุ่นโดยใช้ทวีตเตอร์ ว่าแต่ละแหล่งข่าวนั้นมีความน่าเชื่อถืออยู่ในระดับใดบ้างโดยวัดจากจำนวนของทวีตที่มีความน่าเชื่อถือว่ามีเป็นสัดส่วนเท่าใดของจำนวนทวีตที่มาจากแหล่งข่าวดังกล่าว (เวสเทอร์แมน และคณะ, 2012) พบว่าจำนวนจำนวนผู้ติดตาม (Followers) ในทวีตเตอร์ไม่มีความสัมพันธ์แบบเชิงเส้นต่อการตัดสินคุณวุฒิและความน่าเชื่อถือของบุคคลที่ใช้งานทวีตเตอร์ และพบว่าถ้าจำนวนผู้ติดตามและคนที่ติดตาม (Followings) มีปริมาณใกล้เคียงกันจะมีส่วนช่วยให้ถูกพิจารณาว่าบุคคลดังกล่าวมีคุณวุฒิ แต่หากมีจำนวนผู้ติดตามจำนวนมากและมีคนที่ติดตามจำนวนน้อยนั้นไม่ได้แสดงถึงความเป็นมืออาชีพหรือเป็นผู้เชี่ยวชาญ

ส่วนการตัดสินภาพลักษณ์ของแหล่งข่าวพบว่ามาจากความสามารถ และความน่าเชื่อถือของบุคคล ที่วัดจากเนื้อหาสาระที่ผู้ใช้งานแบ่งปันต่อผู้อื่นมากกว่าการตัดสิน จากจำนวนผู้ติดตามหรืออัตราส่วน ระหว่างผู้ติดตามและคนที่ติดตาม อย่างไรก็ตามงานวิจัยที่กล่าวมาข้างต้นทั้ง 3 งานวิจัยนั้นล้วนแต่ ยังไม่ได้นำเสนอวิธีการสร้างโมเดลวัดความน่าเชื่อถือของทวีต (อัลคาลิฟา และคณะ, 2011) พบว่า การประเมินความน่าเชื่อถือของทวีตด้วยการคำนวณค่า ความเหมือนของเอกสาร (Similarity) เทียบกับชุดเอกสารที่ได้รับการรับรองว่าน่าเชื่อถือนั้นให้ผลที่แม่นยำกว่าผลประเมินความน่าเชื่อถือ จากสมการที่ได้นำเสนอคือ

$$\text{Credibility score} = 0.6S + 0.2IW + 0.1L + 0.1A \quad (1)$$

โดย S คือ ค่าความเหมือนกันของเนื้อหาข้อความ IW คือ จำนวนคำที่เป็นคำภาษาพูด คำที่ไม่เป็นทางการ คำหยาบคาย L คือ จำนวนลิงก์ที่อ้างอิงไปยังแหล่งข้อมูลที่เชื่อถือได้ และ A คือ ค่าคะแนนคุณภาพของบัญชีผู้ใช้ที่เว็บไซต์ TwitterGrader เปิดให้บริการคำนวณค่าให้ และยังได้นำเสนอการคำนวณค่าที่ใช้เป็นจุดแบ่งระดับความน่าเชื่อถือนั้นออกเป็น 3 ช่วงได้แก่ ระดับต่ำ ระดับเฉลี่ย และระดับสูง แต่อย่างไรก็ตามงานวิจัยนี้ใช้ผู้เชี่ยวชาญในการพิจารณาความน่าเชื่อถือของชุด เฉลยเพียง 3 คน และยังไม่ได้พิจารณาค่าจากคุณลักษณะอื่น ๆ ที่สำคัญเช่น แฮชแท็ก จำนวนทวีต และสัญลักษณ์หรือรูปไอคอนแสดงอารมณ์ต่าง ๆ เป็นต้น (แคสทิลโล และคณะ, 2011) นำเสนอการ ประเมินความน่าเชื่อถือของข่าวจากทวีตเตอร์แบบอัตโนมัติด้วยการใช้แนวคิดของต้นไม้ตัดสินใจ แบบ J48 (J48 decision tree) ที่พบว่าให้ความถูกต้องแม่นยำมากที่สุด ด้วยคุณลักษณะเด่นที่ดีที่สุด 15 คุณลักษณะ ระดับความแม่นยำถูกต้องของโมเดลในการจำแนกข้อความข่าวจากทวีตเตอร์ ว่าน่าเชื่อถือหรือไม่นั้นอยู่ในช่วงร้อยละ 70 ถึงร้อยละ 80 โดยการใช้แบบสำรวจกับตัวอย่างโพสต์ ของทวีตเตอร์ (แคสทิลโล และคณะ, 2013) ได้นำเสนอการทดลองสร้างโมเดลจำแนกข้อมูลโดย เลือกใช้วิธีการแรน-ดอมฟอเรสต์ (Random forest) ลอจิสติกส์ (Logistic) และเมตาเลิร์นนิง (Meta learning) ซึ่งเป็นที่นิยมในอันดับต้น ๆ โดยใช้คุณลักษณะที่คัดเลือก 16 คุณลักษณะ พบว่า โมเดลจำแนกข้อมูลที่ใช้แนวคิด ลอจิสติกส์ (Logistic) ให้ผลในการคัดแยกข้อความทวีตที่เป็นข่าวและ ข้อความที่ไม่ใช่ข่าวได้ถูกต้องร้อยละ 78 และสามารถคัดแยกข้อความทวีตที่น่าเชื่อถือกับข้อความ ทวีตที่ไม่น่าเชื่อถือได้ถูกต้องร้อยละ 86 (เซียและคณะ, 2012) นำเสนอการทดลองใช้แนวคิดของเบย์เซียนเน็ตเวิร์ค (Bayesian network) ในการสร้างโมเดลวิเคราะห์ความน่าเชื่อถือของข่าว ที่พบว่าให้ ความถูกต้องแม่นยำในระดับต้น ๆ โดยเปรียบเทียบกับวิธีทางอื่น ๆ เช่น ต้นไม้ตัดสินใจ (J48Decision tree) ซัพพอร์ตเวกเตอร์แมชชีน และ เคทู (K2) อย่างไรก็ตามงานวิจัยที่กล่าวข้างต้น ทั้ง 4 งานวิจัย พบว่ายังไม่มียานวิจัยใดที่พัฒนาเครื่องมือประเมินความน่าเชื่อถือโดยอัตโนมัติที่สามารถ ใช้กับทวีตเตอร์แบบเรียลไทม์ (กัปตา เอ และคณะ, 2014) นำเสนอการใช้แนวคิดของการจัดอันดับ ด้วยซัพพอร์ตเวกเตอร์แมชชีน - ซีน (Rank-SVM) ร่วมกับคุณลักษณะที่คัดเลือก 45 คุณลักษณะในการ สร้างโมเดลจำแนกข้อมูล และนำไปใช้ในระบบประเมินความน่าเชื่อถือที่สร้างขึ้นที่สามารถประเมิน ความน่าเชื่อถือของทวีตเตอร์ได้ 7 ระดับ แบบเรียลไทม์ และร้อยละ 84 ของทวีตมีคะแนนความ น่าเชื่อถือที่ถูกนำไปแสดงแก่ผู้ใช้งานได้ภายใน 6 วินาที แต่อย่างไรก็ตามในงานวิจัยดังกล่าว

ไม่มีการประเมินความแม่นยำของระดับคะแนนความน่าเชื่อถือโดยผู้เชี่ยวชาญ มีแต่การรายงานว่าร้อยละ 43 ของผู้ใช้เห็นด้วยกับคะแนนความน่าเชื่อถือที่นำเสนอโดยระบบ (กัปตาเอเอ็ม และคณะ, 2012) นำเสนอการประเมินความน่าเชื่อถือของเหตุการณ์ในทวีตเตอร์โดยใช้แนวคิดของ อัลกอริทึมจัดอันดับเพจ (PageRank algorithm) โดยพิจารณาคุณสมบัติจากข้อความทวีต ผู้ใช้ทวีตเตอร์ และคุณลักษณะของเหตุการณ์ จากนั้นปรับปรุงคะแนนความน่าเชื่อถือด้วยเทคนิคการเพิ่มประสิทธิภาพด้วยกราฟเหตุการณ์ (Event graph-based optimization) ซึ่งพบว่าให้ความแม่นยำถูกต้องประมาณร้อยละ 86 ซึ่งมากกว่าวิธีการจำแนกโดยใช้ต้นไม้ตัดสินใจ (Decision tree classifier) ที่มีความถูกต้องแม่นยำอยู่ที่ร้อยละ 72 (กัปตา เอ และคณะ, 2012) นำเสนอคุณสมบัติเด่นของเนื้อหา (Content based featured) และ คุณลักษณะเด่นของที่มาของข้อมูล (Source based features) แล้วนำมาใช้ร่วมกับแนวคิดของการจัดอันดับด้วยซัพพอร์ตเวกเตอร์แมชชีนเพื่อทำการจัดลำดับของทวีตตามคะแนนความน่าเชื่อถือในรอบแรก จากนั้นนำผลที่ได้มาจัดลำดับโดยใช้เทคนิคซูโดรีเลแวนซ์ฟีดแบค (Pseudo relevance feedback) แต่งานวิจัยดังกล่าวพบว่าไม่ได้มีการพัฒนาระบบประเมินความน่าเชื่อถือของทวีตและประเมินประสิทธิภาพของระบบ

จากวิธีการของงานวิจัยที่เกี่ยวข้องที่กล่าวมาข้างต้น พบว่ามีเพียงงานวิจัยของ กัปตา เอ และคณะในปี 2014 ที่ได้นำแนวคิดของโมเดลจำแนกข้อมูลมาพัฒนาเป็นระบบประเมินความน่าเชื่อถือข่าวจากทวีตให้สามารถใช้งานในรูปแบบเรียลไทม์ได้ แต่ระบบดังกล่าวให้ผลถูกต้องประมาณร้อยละ 66.40 กับชุดข่าวทวีตที่เป็นภาษาไทยที่ใช้ในการทดลอง งานวิจัยที่นำเสนอในบทความนี้มีวัตถุประสงค์เพื่อนำเสนอและพัฒนาระบบประเมินความน่าเชื่อถือข่าวจากทวีตที่เป็นภาษาไทยที่มีความถูกต้องและมีประสิทธิภาพมากยิ่งขึ้น

4. วิธีการวิจัย

4.1. การเก็บข้อมูล

ผู้วิจัยเก็บข้อมูลทวีตผ่านเอพีไอของทวีตเตอร์โดยใช้บริการ GET search/tweets โดยทำการเก็บข้อมูลทวีตในหัวข้อเด็กหาย หรือคนหายจำนวน 158,996 ทวีต เพื่อใช้สำหรับสร้างชุดข้อมูลเทรนนิ่งแล้วนำมาใช้ในการพัฒนาโมเดลจำแนกข้อมูลที่ใช้ประเมินความน่าเชื่อถือของข่าวจากทวีตเตอร์ และเก็บข้อมูลทวีตในหัวข้ออุบัติเหตุจำนวน 4,962 ทวีต เพื่อใช้สร้างชุดข้อมูลทดสอบที่ใช้วัดประสิทธิภาพของโมเดลจำแนกข้อมูลที่สร้างขึ้นเปรียบเทียบกับระบบ Tweetcred ของงานวิจัย กัปตา เอ และคณะ

4.2. การสร้างชุดข้อมูล (Data set) จากข้อมูลทวีต

ผู้วิจัยใช้ทวีตจำนวน 1,000 ทวีตที่สุ่มเลือกจากข้อมูลทวีตในหัวข้อเด็กหาย หรือ คนหาย และทวีตจำนวน 500 ทวีตที่สุ่มเลือกจากข้อมูลทวีตในหัวข้ออุบัติเหตุ โดยทวีตที่สุ่มเลือกมาทั้ง 2 ชุดนั้นจะถูกนำมากำหนดความน่าเชื่อถือโดยผู้เชี่ยวชาญจำนวน 5 คนช่วยกันประเมินความน่าเชื่อถือของแต่ละทวีตผ่านแบบสอบถามออนไลน์ที่พัฒนาขึ้น ที่มีข้อความข่าวจากทวีต ชื่อเจ้าของทวีต จำนวนรีทวีตประกอบการตัดสินใจ และมีระดับความน่าเชื่อถือของทวีตให้เลือกจากน้อยไปมากคือระดับ 1 ถึงระดับ 5 ให้เลือก ส่วนคุณลักษณะที่ใช้ในการสร้างชุดข้อมูลเทรนนิ่งเพื่อใช้ในการสร้างโมเดลจำแนกข้อมูลตามระดับความน่าเชื่อถือได้แก่

1. สัญลักษณ์ที่แสดงอารมณ์ในข้อความทวิต เช่น :) (^_^
2. คำที่เป็นภาษาพูดหรือคำหยาบคายในข้อความทวิต เช่น มึง กู
3. แฮชแท็กที่ใช้เพื่อความสนุกสนานในข้อความทวิต เช่น #ต้ง
4. ไอคอนแสดงอารมณ์หรือตัวการ์ตูนในข้อความทวิต
5. การกล่าวถึงผู้ใช้ที่เป็นนักข่าวหรือผู้ใช้ที่เป็นหน่วยงานที่ให้บริการข่าวสารในข้อความทวิต เช่น @js100radio
6. ลิงก์ของเว็บไซต์ที่ให้ข้อมูลเพิ่มเติมในข้อความทวิต
7. ลิงก์ของเว็บไซต์ที่ให้ข้อมูลเพิ่มเติมนั้นมาจากเว็บไซต์ข่าว
8. รูปภาพประกอบถูกแนบมาแสดงในข้อความทวิต
9. รูปภาพประกอบที่แสดงในข้อความทวิต นั้นมาจากเว็บไซต์ข่าว
10. ข้อความทวิตนั้นถูกทวิตหรือรีทวิตโดยผู้ใช้ที่เป็นนักข่าวหรือผู้ใช้ที่เป็นหน่วยงานที่ให้บริการข่าวสาร เช่น @js100radio

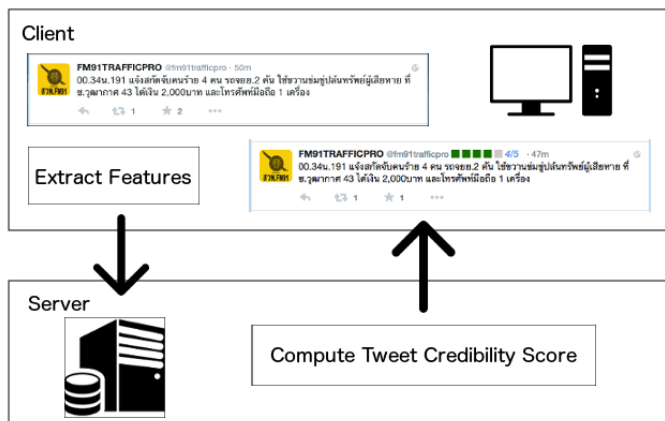
คุณลักษณะ 10 คุณลักษณะที่เลือกมาใช้ในการทดลองสร้างโมเดลจำแนกข้อมูลนี้พิจารณาจากงานวิจัย (กัปดา เอ และคณะ, 2014) (มอริส และคณะ, 2012) (แคสทิลโล และคณะ, 2011) ที่พบว่าถูกใช้ร่วมกัน และเลือกคุณลักษณะที่เกี่ยวข้องกับภาษา หรือแหล่งข้อมูลที่มีเฉพาะในท้องถิ่น เช่น คุณลักษณะที่ได้จากการตรวจสอบว่าพบคำที่เป็นภาษาพูดหรือคำหยาบคายในข้อความทวิตหรือไม่ หรือ คุณลักษณะที่ได้จากการตรวจสอบว่าพบแฮชแท็กที่ใช้เพื่อความสนุกสนานในข้อความทวิตหรือไม่ เป็นต้น โดยค่าจากคุณลักษณะที่นำมาพิจารณานี้ จะอยู่ในรูปแบบของค่าความจริง คือเป็นจริง หรือ เป็นเท็จ

4.3. การสร้างโมเดลจำแนกข้อมูล

อัลกอริทึมที่ใช้ในการสร้างโมเดลจำแนกข้อมูลสำหรับใช้ประเมินความน่าเชื่อถือของข่าวจากทวิตเตอร์ที่เลือกมาทดลองประกอบไปด้วยซัพพอร์ตเวกเตอร์แมชชีน (Support vector machine) ซึ่งเป็นอัลกอริทึมหนึ่งที่นิยมทางด้านแมชชีนเลิร์นนิง (Machine learning) ที่ใช้ในงานจำแนกข้อมูลออกเป็นกลุ่ม โดยใช้เครื่องมือที่ชื่อลิปเอสวีเอ็ม (LibSVM) (ซาง และลิน, 2011) และอัลกอริทึมออเดอร์ลอจิสติกส์ (Ordered logistic) ที่เป็นอัลกอริทึมหนึ่งที่นิยมในกลุ่มนักสถิติใช้สร้างโมเดลในการวิเคราะห์และจำแนกข้อมูลออกเป็นกลุ่มโดยใช้เครื่องมือที่ชื่ออาร์ (R) โดยใช้แพ็คเกจที่ชื่อแมส (MASS) (เวเนเบิลซ์ และลิปส์, 2002) เนื่องจาก 2 อัลกอริทึมนี้สามารถจำแนกข้อมูลได้มากกว่า 2 กลุ่ม ซึ่งตรงกับความต้องการของระบบประเมินความน่าเชื่อถือที่ผู้วิจัยนำเสนอต้องการให้สามารถคำนวณและแสดงผลระดับความน่าเชื่อถือแก่ผู้ใช้งานได้ 5 ระดับ

4.4. การพัฒนาระบบประเมินความน่าเชื่อถือข่าวจากทวิตเตอร์

ระบบจะประกอบไปด้วยโปรแกรมที่ทำงานร่วมกัน 2 ส่วนคือ โปรแกรมที่ทำงานบนฝั่งผู้ใช้งานที่เป็นส่วนขยายของโครม (Chrome extension) ที่มีชื่อว่า Thaitweetcred และโปรแกรมที่ทำงานบนเซิร์ฟเวอร์ดังภาพที่ 1 ที่ทำหน้าที่อ่านข้อมูลคุณสมบัติจากทวิตของผู้ใช้งานแล้ว ส่งค่ามาประมวลผลเป็นระดับความน่าเชื่อถือโดยอัตโนมัติซึ่งแสดงแก่ผู้ใช้งานแบบเรียลไทม์ดังภาพที่ 2



ภาพที่ 1 แนวคิดการทำงานของระบบประเมินความน่าเชื่อถือที่นำเสนอ



รูปที่ 2 การแสดงระดับความน่าเชื่อถือของทวีตของระบบที่นำเสนอ

4.5. การวัดประสิทธิภาพของระบบประเมินความน่าเชื่อถือข่าว

จากทวีตเตอร์ที่นำเสนอเทียบกับระบบของ Tweetcred มีขั้นตอนดังต่อไปนี้คือนำข้อความ 500 ทวีตที่เป็นชุดเดียวกันกับชุดทดสอบนั้นไปทดสอบให้ระบบประเมินความน่าเชื่อถือ ที่นำเสนอ ประเมินความน่าเชื่อถือออกมาให้ โดยระบบจะมีรูปแบบการแสดงระดับความน่าเชื่อถือเป็น 5 ระดับ และนำข้อความ 500 ทวีต ที่เป็นชุดเดียวกันกับชุดทดสอบนั้น ไปทดสอบให้ระบบของ Tweetcred ประเมินความน่าเชื่อถือออกมาให้ โดยระบบจะมีรูปแบบการแสดงระดับความน่าเชื่อถือเป็น 7 ระดับ แล้วปรับรูปแบบของระดับความน่าเชื่อถือของผลที่ได้ รวมทั้งผลในชุดเจลยนั้นให้มีรูปแบบของระดับความน่าเชื่อถือที่ตรงกัน โดยเราทำการเปลี่ยนค่าระดับความน่าเชื่อถือ จากผลทั้ง 3 ชุดนั้น ให้มีค่าระดับความน่าเชื่อถือเหลือเพียง 2 ระดับ คือ ระดับ 0 ไม่น่าเชื่อถือกับ ระดับ 1 น่าเชื่อถือ โดยผลที่มีระดับความน่าเชื่อถือในรูปแบบ 7 ระดับนั้น หากมีระดับความน่าเชื่อถืออยู่ในช่วง 1-4 ให้เปลี่ยนเป็นระดับ 0 คือไม่น่าเชื่อถือ และหากมีระดับความน่าเชื่อถืออยู่ในช่วง 5-7 ให้เปลี่ยนเป็น ระดับ 1 คือ น่าเชื่อถือ ส่วนผลที่มีระดับความน่าเชื่อถือในรูปแบบ 5 ระดับนั้น หากมีระดับความน่าเชื่อถืออยู่ในช่วง 1-3 ให้เปลี่ยนเป็น ระดับ 0 คือไม่น่าเชื่อถือ และหากมีระดับความน่าเชื่อถืออยู่ในช่วง 4-5

ให้เปลี่ยนเป็น ระดับ 1 คือนำเชื่อถือ แล้วนำผลที่ได้นั้นไปตรวจสอบว่าแต่ละทวิตมีระดับความน่าเชื่อถือที่ตรงกับผลในชุดเฉลี่ยหรือไม่ และวัดประสิทธิภาพของโมเดลจำแนกข้อมูลด้วยการหาค่าความถูกต้อง (Accuracy) ค่าความแม่นยำ (Precision) ค่ารีคอล (Recall) และค่าเอฟเมซเซอร์ (F-Measure)

5. ผลการวิจัย

5.1 ผลการทดลองเพื่อเลือกอัลกอริทึมที่เหมาะสมสำหรับใช้สร้างโมเดลจำแนกข้อมูล

จากตารางที่ 1 และ 2 แสดงผลการวัดประสิทธิภาพจากโมเดลจำแนกข้อมูลทั้ง 2 อัลกอริทึมด้วยวิธีการแบ่งข้อมูลชุดทดสอบเป็น 5 ส่วน แล้วนำมาหาค่าเฉลี่ย พบว่าค่าเฉลี่ยของความถูกต้องในการจำแนกข้อมูลระดับความน่าเชื่อถือจากอัลกอริทึมซัพพอร์ตเวกเตอร์แมชชีนจะอยู่ที่ร้อยละ 71.96 ส่วนอัลกอริทึมอเดอร์ลอจิสติกส์ นั้นจะอยู่ที่ร้อยละ 71.18 ซึ่งแตกต่างกันเพียงร้อยละ 0.78 และเมื่อพิจารณาประสิทธิภาพความถูกต้องในการจำแนกข้อมูลของแต่ละคลาสจากค่าเฉลี่ยของค่าเอฟเมซเซอร์จะพบว่าค่าที่ได้จากทั้ง 2 อัลกอริทึมนั้นเป็นไปในทางเดียวกันคือ โมเดลจำแนกข้อมูลจะสามารถจำแนกข้อมูลได้ถูกต้องมากที่สุดคือคลาสที่มีระดับความน่าเชื่อถือระดับ 1 ส่วนคลาสที่มีระดับความน่าเชื่อถือระดับ 5 และ 3 นั้นจะเป็นคลาสที่จำแนกข้อมูลได้ถูกต้องรองลงมาตามลำดับ ส่วนในคลาสที่มีระดับความน่าเชื่อถือระดับ 2 และ 4 นั้นจะพบว่าค่าเฉลี่ยของค่าเอฟเมซเซอร์ที่ได้ค่อนข้างต่ำมากคือมีค่าอยู่ในช่วงที่ไม่เกินร้อยละ 10 ซึ่งเป็นผลจากการเลือกระดับความน่าเชื่อถือให้แต่ละทวิตจากผู้เชี่ยวชาญนั้นเป็นความเห็นของผู้เชี่ยวชาญแต่ละคนที่ไม่ได้ทำการตกลงเกณฑ์การพิจารณาของความน่าเชื่อถือแต่ละระดับที่ตรงกันก่อนอีกด้วย และหากพิจารณาค่าความคลาดเคลื่อนจากค่ารูทมีนสแควร์ (Root mean square) จะพบว่าคลาดเคลื่อนอยู่ไม่เกิน 2 ระดับ ผู้วิจัยเลือกอัลกอริทึมซัพพอร์ตเวกเตอร์ – แมชชีน ในการสร้างโมเดลจำแนกข้อมูลเพื่อนำไปใช้ ในระบบประเมินความน่าเชื่อถือเนื่องจากให้ความถูกต้องในการจำแนกที่สูงกว่าอยู่ร้อยละ 0.78 และมีเครื่องมือที่อินเตอร์เฟสที่นำมาพัฒนาระบบสะดวกกว่าอัลกอริทึมอเดอร์ลอจิสติกส์

5.2 ผลการวิเคราะห์ค่าความถี่ของคุณลักษณะที่ใช้จำแนกข้อมูลจากชุดข้อมูลเทรนนิ่ง

จากตารางที่ 3 พบว่าคุณลักษณะที่ 1 ถึง คุณลักษณะที่ 4 เป็นคุณสมบัติด้านลบ จะพบว่ามีโอกาสพบคุณสมบัติดังกล่าวมากในทวิตที่มีระดับความน่าเชื่อถือน้อย และมีโอกาสพบน้อยลงในทวิตที่มีระดับความน่าเชื่อถือมาก ในทำนองเดียวกันคุณลักษณะที่ 5 ถึง 10 เป็นคุณสมบัติด้านบวก จะพบว่ามีโอกาสพบคุณสมบัติดังกล่าวน้อยในทวิตที่มีระดับความเชื่อถือน้อย และมีโอกาสพบมากขึ้นในทวิตที่มีระดับความน่าเชื่อถือมาก จากการวิเคราะห์ค่าความน่าจะเป็นที่ใช้ตรวจสอบสมมติฐาน (P-Value) ว่าสัมประสิทธิ์ของตัวแปรของคุณลักษณะที่ใช้ในการจำแนกข้อมูลนั้นเป็น 0 หรือไม่ของแต่ละคุณลักษณะจากอัลกอริทึมอเดอร์ลอจิสติกส์จะพบว่า คุณลักษณะที่ 4 และ 7 นั้นมีค่าดังกล่าวไม่น้อยกว่า 0.05 ซึ่งแปลผลได้ว่าคุณลักษณะทั้ง 2 นี้ไม่ส่งผลต่อการจำแนกข้อมูลของโมเดล และเมื่อทดลองตัดคุณลักษณะทั้ง 2 นี้ออกแล้วทำการสร้างและวัดประสิทธิภาพของโมเดลจำแนกข้อมูลอีกรอบ พบว่าจำนวนทวิตที่ระบบประเมินความน่าเชื่อถือได้ตรงกับผู้เชี่ยวชาญนั้นแตกต่างจากเดิมเพียง 1 ทวิต ซึ่งเป็นผลมาจากฐานข้อมูลที่ใช้ตรวจสอบคุณสมบัติทั้ง 2 นั้นยังมีข้อมูล

น้อย และยังไม่ครอบคลุม จึงทำให้คุณลักษณะทั้ง 2 นั้นยังไม่ส่งผลและช่วยทำให้การจำแนกข้อมูลตามความน่าเชื่อถือได้ถูกต้องแม่นยำ

5.3 ผลการวัดประสิทธิภาพของระบบประเมินความน่าเชื่อถือที่นำเสนอกับระบบของ Tweetcred

ผู้วิจัยเลือกใช้ข้อมูลทวีตในหัวข้ออุบัติเหตุที่เป็นข้อมูลคนละกลุ่มกับชุดข้อมูลที่ใช้สร้างโมเดลจำแนกข้อมูล โดยทำการสุ่มข้อมูลทวีตมาจำนวน 500 ทวีต แล้วแบ่งข้อมูลชุดทดสอบนั้นออกเป็น 5 ส่วน แล้วทำการทดสอบโดยการให้ระบบประเมินความน่าเชื่อถือของข่าวจากทวีตเตอร์ทั้ง 2 ระบบนั้นประเมินความน่าเชื่อถือของแต่ละทวีตแล้วนำมาเทียบกับระดับความน่าเชื่อถือได้จากผู้เชี่ยวชาญจะพบว่าความถูกต้องในการจำแนกข้อมูลโดยใช้ระบบประเมินความน่าเชื่อถือที่นำเสนอมีค่าเฉลี่ยอยู่ที่ร้อยละ 65 ส่วนค่าร้อยละความถูกต้องในการจำแนกข้อมูลของระบบ Tweetcred จากข้อมูลชุดทดสอบเดียวกันแสดงในตารางที่ 4 นั้นพบว่ามีความเฉลี่ยอยู่ที่ร้อยละ 66.40 จากนั้นผู้วิจัยได้ทำการวิเคราะห์ข้อมูลทวีตที่ใช้เป็นชุดทดสอบเพิ่มเติมพบว่ามีทวีตที่มาจากผู้ใช้งานทวีตเตอร์ ที่เป็นผู้สื่อข่าวหรือหน่วยงานที่แจ้งข่าวเกี่ยวกับหัวข้ออุบัติเหตุอยู่เป็นจำนวนมากและบัญชีผู้ใช้เหล่านี้ยังไม่ได้เก็บลงฐานข้อมูลที่ใช้ตรวจสอบคุณลักษณะที่ 10 คือ ข้อความทวีตนั้นถูกทวีตหรือรีทวีตโดยผู้ใช้ที่เป็นนักข่าวหรือผู้ใช้ที่เป็นหน่วยงานที่ให้บริการข่าวสาร หลังจากเพิ่มบัญชีผู้ใช้งานทวีตเตอร์ที่เป็นผู้สื่อข่าวหรือหน่วยงานที่หาหน้าที่ให้ข่าวเกี่ยวกับหัวข้ออุบัติเหตุได้แก่ motorway_th, Ruamduay, traffy, CALL192, police7oc, highwaypolice80, GURUJARAJORN, longdotraffic จากนั้นทำการวัดความถูกต้องในการจำแนกข้อมูลของระบบประเมินความน่าเชื่อถือที่นำเสนอใหม่อีกรอบจะได้ผลลัพธ์ดังตารางที่ 5 ที่พบว่าร้อยละความถูกต้องในการจำแนกข้อมูลนั้นมีค่าเป็นร้อยละ 90.60

5.4 ผลเปรียบเทียบประสิทธิภาพด้านเวลาที่ใช้ของระบบประเมินความน่าเชื่อถือที่นำเสนอ กับระบบของ Tweetcred

ผู้วิจัยได้ทำการบันทึกค่าเวลาที่ใช้ในการส่งส่งเฮชทีทีพีรีควีสตั้งแต่การส่งข้อมูลคุณลักษณะไปประมวลผลที่เซิร์ฟเวอร์จนกระทั่งส่งผลลัพธ์กลับมาแสดงแก่ผู้ใช้งาน จำนวน 5 ชุดจากระบบประเมินความน่าเชื่อถือทั้ง 2 ระบบพบว่าระบบประเมินความน่าเชื่อถือที่นำเสนอใช้เวลาโดยเฉลี่ยอยู่ที่ 184.40 มิลลิวินาที ส่วนระบบของ Tweetcred จะใช้เวลาโดยเฉลี่ยอยู่ที่ 1,912.67 มิลลิวินาที ซึ่งเป็นผลมาจากที่ระบบที่นำเสนอนั้นใช้จำนวนคุณลักษณะ 10 คุณลักษณะที่สามารถอ่านค่าได้โดยตรงจากโปรแกรมที่นำเสนอในฝั่งผู้ใช้ ซึ่งก็คือโครมเอกซ์เทนชันของ Thaitweetcred แล้วนำไปประมวลผลเพื่อให้ได้ระดับคะแนนความน่าเชื่อถือ ในขณะที่ระบบ Tweetcred ใช้จำนวนคุณลักษณะ 45 คุณลักษณะ ที่อ่านค่ามาจากทวีตเตอร์เอพีไอ จึงทำให้เวลาที่ใช้ในการประมวลผลของระบบที่นำเสนอนั้นลดลงถึงร้อยละ 90.35

5.5 ความพึงพอใจของผู้ใช้งานต่อระบบประเมินความน่าเชื่อถือที่นำเสนอ

พบว่าเมื่อผู้เข้ามาใช้งานระบบทั้งสิ้น 45 คนและมีข้อมูลป้อนกลับที่เป็นความเห็นจากผู้ใช้งานจำนวน 365 ความเห็น เห็นด้วยกับระดับความน่าเชื่อถือที่ได้จากระบบประเมินความน่าเชื่อถือมีจำนวน 241 ความเห็น ซึ่งคิดเป็นร้อยละ 66.03 และไม่เห็นด้วยกับระดับความน่าเชื่อถือที่ได้จากระบบประเมินความน่าเชื่อถือจำนวน 124 ความเห็น คิดเป็นร้อยละ 33.97 ซึ่งน่าจะเป็นผล

มาจากข่าวที่ผู้ใช้งานทวิตเตอร์ทดลองใช้นั้นเป็นข่าวที่มาจากหลากหลายประเภทข่าว รูปแบบการจำแนกข้อมูลจากหัวข้อเด็กหาย อาจจะไม่สามารถนำไปใช้จำแนกข้อมูลข่าวตามความน่าเชื่อถือในหัวข้ออื่นหรืออีกประเภทอื่นได้ดีเท่าที่ควร

6. สรุปและอภิปรายผล

งานวิจัยนี้ได้ศึกษาเพื่อหาวิธีการประเมินความน่าเชื่อถือของข่าวจากทวิตเตอร์ที่เป็นภาษาไทยให้สามารถแสดงระดับความน่าเชื่อถือได้ 5 ระดับ โดยใช้โมเดลจำแนกข้อมูลมาจำแนกข้อความทวิตตามระดับความน่าเชื่อถือ พบว่าโมเดลจำแนกข้อมูลที่สร้างจากอัลกอริทึมซัพพอร์ตเวกเตอร์แมชชีนจะมีความแม่นยำสูงที่สุดเมื่อเทียบกับอัลกอริทึมอเดอร์ลอจิสติกส์ โดยใช้คุณลักษณะในการจำแนกข้อมูล 10 คุณลักษณะเหมือนกัน และได้นำโมเดลจำแนกข้อมูลมาพัฒนาระบบประเมินความน่าเชื่อถือแบบเรียลไทม์ให้ผู้ใช้งานทวิตเตอร์นั้นสามารถนำไปใช้ประเมินความน่าเชื่อถือของข่าวในทวิตเตอร์เบื้องต้นได้ แนวทางในการพัฒนาโมเดลจำแนกข้อมูลเพื่อประเมินความน่าเชื่อถือให้มีความแม่นยำมากยิ่งขึ้น อาจทำได้โดยการ เพิ่มข้อมูลที่เป็นชุดข้อมูลเทรนนิ่งในหัวข้ออื่นๆ ที่ผู้ใช้ให้ความสำคัญในเรื่องความน่าเชื่อถือ เช่น อุบัติเหตุ ภัยพิบัติ เหตุฉวน เหตุร้าย หรือหัวข้อเฉพาะที่ต้องการนำโมเดลจำแนกข้อมูลไปใช้งาน เพิ่มคุณลักษณะของการพิจารณาความน่าเชื่อถือของทวิตที่เกี่ยวข้อง เช่น ทวิตของคนที่ยืนยันตัวตน เพิ่มคุณลักษณะของการพิจารณาความน่าเชื่อถือจากกลุ่มบุคคลที่ทำการรีทวีตข้อความข่าว เพิ่มคุณลักษณะการพิจารณาความน่าเชื่อถือจากการตรวจสอบว่าในทวิต มีการระบุหรืออ้างอิงถึง ชื่อบุคคล เบอร์โทร หรือสถานที่หรือไม่ และเพิ่มบัญชีของนักข่าวในฐานะข้อมูลของระบบโดยอัตโนมัติจากการแนะนำของบุคคลที่น่าเชื่อถือและผู้ใช้งานจำนวนมาก

ตารางที่ 1 ผลการวัดประสิทธิภาพของโมเดลจำแนกข้อมูลที่ใช้อัลกอริทึมซัพพอร์ตเวกเตอร์แมชชีน

คลาส	จำนวนที่ จำแนก โดย ผู้เชี่ยวชาญ	จำนวนที่ จำแนก โดย โมเดล	จำนวนที่ โมเดล จำแนกได้ ถูกต้อง	Accuracy (%)	Precision (%)	Recall (%)	F- Measure (%)	Root mean square error
ความน่าเชื่อถือ ระดับ 1	587.20	631	550.20	71.96	87.23	93.78	90.28	0.48
ความน่าเชื่อถือ ระดับ 2	121.80	6.20	0.80		15.33	0.66	1.21	1.05
ความน่าเชื่อถือ ระดับ 3	123.60	280	106.60		38.37	85.53	51.79	0.72
ความน่าเชื่อถือ ระดับ 4	75.00	7.80	3.80		54.26	5.78	9.89	1.18
ความน่าเชื่อถือ ระดับ 5	92.40	75.00	58.2		77.60	67.38	70.84	1.16

ตารางที่ 2 ผลการวัดประสิทธิภาพของโมเดลจำแนกข้อมูลที่ใช้อัลกอริทึมออเดอร์ลอจิสติกส์

คลาส	จำนวนที่ จำแนก โดย ผู้เชี่ยวชาญ	จำนวน ที่ จำแนก โดย โมเดล	จำนวน ที่โมเดล จำแนก ได้ ถูกต้อง	Accuracy (%)	Precision (%)	Recall (%)	F- Measure (%)	Root mean square error
ความน่าเชื่อถือ ระดับ 1	587.20	636	553.20	71.18	87.02	94.31	90.40	0.46
ความน่าเชื่อถือ ระดับ 2	121.80	7.20	0.60		4.29	0.43	0.77	1.13
ความน่าเชื่อถือ ระดับ 3	123.60	266.20	99.40		37.61	80.30	50.06	0.82
ความน่าเชื่อถือ ระดับ 4	75.00	8.20	1.60		19.74	2.22	3.94	1.21
ความน่าเชื่อถือ ระดับ 5	92.40	82.40	57.00		69.32	65.97	66.03	1.18

ตารางที่ 3 ความถี่ของแต่ละคุณลักษณะที่ใช้จำแนกข้อมูลที่พบในแต่ละคลาส

ระดับ ความ น่าเชื่อถือ	จำนวน	ความถี่ของคุณลักษณะที่พบในข้อความที่วัดจากชุดข้อมูลเทรนนิ่งโมเดล									
		จำแนกข้อมูล									
		คุณลักษณะที่									
		1	2	3	4	5	6	7	8	9	10
ระดับ 1	2,936	637	1,145	201	250	6	112	0	174	0	3
ระดับ 2	609	55	162	21	45	10	166	9	178	6	7
ระดับ 3	618	11	48	4	29	49	280	9	299	6	11
ระดับ 4	375	8	12	4	23	52	216	21	148	37	63
ระดับ 5	462	4	13	0	3	103	262	51	216	171	291

ตารางที่ 4 ผลการวัดประสิทธิภาพของระบบประเมินความน่าเชื่อถือของ Tweetcred

คลาส	จำนวนที่ จำแนก โดย ผู้เชี่ยวชาญ	จำนวน ที่ จำแนก โดย โมเดล	จำนวน ที่โมเดล จำแนก ได้ ถูกต้อง	Accuracy (%)	Precision (%)	Recall (%)	F- Measure (%)	Root mean square error
น่าเชื่อถือ	59.80	33.80	30.00		88.62	50.12	63.75	0.31
ไม่ น่าเชื่อถือ	40.20	66.20	36.40	66.40	55.08	90.45	68.35	0.70

ตารางที่ 5 ผลการวัดประสิทธิภาพของระบบประเมินความน่าเชื่อถือของงานวิจัยที่นำเสนอ

คลาส	จำนวนที่ จำแนก โดย ผู้เชี่ยวชาญ	จำนวน ที่ จำแนก โดย โมเดล	จำนวนที่ โมเดล จำแนกได้ ถูกต้อง	Accuracy (%)	Precision (%)	Recall (%)	F- Measure (%)	Root mean square error
น่าเชื่อถือ	59.80	56.80	53.60		94.51	89.54	91.77	0.28
ไม่ น่าเชื่อถือ	40.20	43.20	37.00	90.60	86.64	91.92	88.86	0.30

7. เอกสารอ้างอิง

- A. Gupta, P. Kumaraguru, C. Castillo, and P. Meier (2014). "TweetCred: A Real- time Web-based System for Assessing Credibility of Content on Twitter," **SocInfo**, pp. 228–243, 2014.
- A. Gupta and P. Kumaraguru (2012). "Credibility ranking of tweets during high impact events," **1st Workshop on Privacy and Security in Online Social Media**, 2012, pp. 2-8.
- C. Castillo, M. Mendoza, and B. Poblete (2011). "Information credibility on twitter," in **20th international conference on World Wide Web**, 2011, pp. 675-684, Hyderabad: ACM.
- C. Castillo, M. Mendoza, and B. Poblete (2013). "Predicting Information Credibility in Time-Sensitive Social Media," **Internet Research**, 23(5), pp. 560-588, 2013.
- C.-C. Chang and C.-J. Lin. (2011). "LIBSVM : a library for support vector machines", in **ACM Transactions on Intelligent Systems and Technology**, 2:27:1--27:27, 2011.

- D. Westerman, P. R. Spence and B. Van Der Heide (2012). "A social network as information: The effect of system generated reports of connectedness on credibility on Twitter", **Computers in Human Behavior**, 28, pp. 199-206, 2012.
- H.S. Al-Khalifa and R.M. Al-Eidan (2011). "An experimental system for measuring the credibility of news content in Twitter," **International Journal of Web Information Systems**, 7,(2), pp. 130-151, 2011.
- J. O'Donovan, B. Kang, G. Meyer, T. Hllerer, and S. Adali (2012). "Credibility in context: an analysis of feature distributions in twitter," **2012 ASE/IEEE International Conference on Social Computing, Social Com**, pp. 167–174, Amsterdam: IEEE.
- M. Gupta, P. Zhao, and J. Han (2012) "Evaluating event credibility on twitter," in **SIAM International Conference on Data Mining (SDM'12)**, 2012, pp. 153-164.
- M.R. Morris, S. Counts, A. Roseway, A. Hoff, and J. Schwarz (2012). "Tweeting is believing?: understanding microblog credibility perceptions," **presented in the ACM 2012 conference on Computer Supported Cooperative Work, February 11-15, 2012, Seattle, Washington, USA.**
- R. Thomson, et al. (2012). "Trusting Tweets: The Fukushima Disaster and Information Source Credibility on Twitter," **9th International Conference on Information Systems for Crisis Response and Management**, 2012, pp. 1-10.
- Twitter, Inc. (2015). **Twitter Usage/Company Facts**. [Online] Available HTTP: <https://about.twitter.com/company>, access on 27/06/2015
- Venables, W. N. & Ripley, B. D. (2002) **Modern Applied Statistics with S**. Fourth Edition. Springer, New York. ISBN 0-387-95457-0.
- X. Xia, X. Yang, C. Wu, S. Li and L. Bao (2012). "Information Credibility on Twitter in Emergency Situation," **2012 Pacific Asia conference on Intelligence and Security Informatics**, 2012, pp. 45-59.