

A PERFORMANCE EVALUATION OF ACTIVE LEARNING BY GREEDY APPROACH AND THOMPSON SAMPLING APPROACH ON TEXT-BASED DATA

Arnon PROMJUN¹ and Seksan KIATSUPAIBUL¹

1 Faculty of Commerce and Accountancy, Chulalongkorn University, Bangkok, Thailand;
arnon11022542@gmail.com

ARTICLE HISTORY

Received: 7 April 2025

Revised: 21 April 2025

Published: 6 May 2025

ABSTRACT

This study investigates the effectiveness of active learning strategies for selecting informative data points to enhance model performance in text classification tasks. Specifically, it compares random sampling, greedy selection, and Thompson sampling with Laplace approximation in the context of labeling tweets related to tourism in Bangkok. A logistic regression model was trained over 100 iterations using data selected by each method. The results indicate that greedy selection consistently outperformed the other approaches in the early stages, enabling rapid model improvement. However, its effectiveness declined in later stages as the availability of informative tweets decreased. Thompson sampling with Laplace approximation exhibited slower initial performance and required more time for data selection, but demonstrated steady improvement across iterations. In contrast, random sampling was the fastest method but failed to significantly enhance model performance, maintaining low accuracy throughout the experiment. These findings suggest that greedy selection is well-suited for applications requiring quick learning, while Thompson sampling holds promise for long-term learning scenarios. The insights gained from this research can inform the development of active learning frameworks in natural language processing tasks, including sentiment analysis and customer opinion mining across various industries.

Keywords: Active Learning, Greedy, Thompson Sampling, Laplace Approximation

CITATION INFORMATION: Promjun, A., & Kiatsupaibul, S. (2025). A Performance Evaluation of Active Learning by Greedy Approach and Thompson Sampling Approach on Text-Based Data. *Procedia of Multidisciplinary Research*, 3(5), 36

การประเมินประสิทธิภาพของการเรียนรู้เชิงรุกโดยวิธีการเลือกแบบละโมบ และวิธีการสุ่มตัวอย่างของทอมป์สันกับข้อมูลรูปแบบข้อความ

อานนท์ พรหมจรรย์¹ และ เสกสรร เกียรติสุไพบูรณ์¹

1 คณะพาณิชยศาสตร์และการบัญชี จุฬาลงกรณ์มหาวิทยาลัย; arnon11022542@gmail.com

บทคัดย่อ

งานวิจัยนี้ศึกษาประสิทธิภาพของวิธีการเรียนรู้แบบแอคทีฟในการคัดเลือกข้อมูลที่มีประโยชน์สูงสุดสำหรับการทำป้ายกำกับข้อมูล โดยเปรียบเทียบสามวิธี ได้แก่ การสุ่มแบบสุ่มทั่วไป, การเลือกแบบละโมบ และวิธีการสุ่มตัวอย่างของทอมป์สันร่วมกับการประมาณค่าแบบลาปลาซ ในบริบทของทวิตที่เกี่ยวข้องกับการท่องเที่ยวในกรุงเทพฯ โดยใช้แบบจำลองการวิเคราะห์ถดถอยโลจิสติกตลอด 100 รอบการคัดเลือกข้อมูล ผลการทดลองพบว่า วิธีการเลือกแบบละโมบให้ประสิทธิภาพสูงที่สุดในช่วงต้นของการเรียนรู้ เนื่องจากสามารถปรับปรุงแบบจำลองได้อย่างรวดเร็ว แต่ประสิทธิภาพลดลงเมื่อจำนวนข้อมูลที่มีสาระลดน้อยลง ส่วนวิธีการสุ่มตัวอย่างของทอมป์สันแม้จะใช้เวลามากกว่าและมีประสิทธิภาพต่ำในช่วงต้น แต่แบบจำลองสามารถพัฒนาได้อย่างต่อเนื่องในระยะยาว ขณะที่การสุ่มแบบทั่วไปแม้จะใช้เวลาน้อยที่สุด แต่ไม่สามารถส่งเสริมการเรียนรู้ของแบบจำลองได้อย่างมีนัยสำคัญ ผลการวิจัยชี้ให้เห็นว่าวิธีการเลือกแบบละโมบเหมาะสมกับสถานการณ์ที่ต้องการผลลัพธ์รวดเร็ว ขณะที่วิธีการสุ่มตัวอย่างของทอมป์สันมีศักยภาพในการเรียนรู้ระยะยาว และสามารถนำไปประยุกต์ใช้ในงานวิเคราะห์ข้อความ เช่น การวิเคราะห์ความรู้สึกของผู้ใช้โซเชียลมีเดีย หรือการประเมินความคิดเห็นของลูกค้าในอุตสาหกรรมต่างๆ

คำสำคัญ: การเรียนรู้เชิงรุก, วิธีการเลือกแบบละโมบ, วิธีการสุ่มตัวอย่างของทอมป์สัน, การประมาณแบบลาปลาซ

ข้อมูลการอ้างอิง: อานนท์ พรหมจรรย์ และ เสกสรร เกียรติสุไพบูรณ์. (2568). การประเมินประสิทธิภาพของการเรียนรู้เชิงรุกโดยวิธีการเลือกแบบละโมบและวิธีการสุ่มตัวอย่างของทอมป์สันกับข้อมูลรูปแบบข้อความ. *Procedia of Multidisciplinary Research*, 3(5), 36

บทนำ

ในปัจจุบัน ข้อมูลมีบทบาทสำคัญในการวิเคราะห์แนวโน้มและพฤติกรรมต่างๆ โดยเฉพาะในยุคที่ข้อมูลสามารถถูกเก็บรวบรวมและประมวลผลได้ในปริมาณมหาศาล การนำข้อมูลเหล่านี้มาวิเคราะห์เพื่อทำนายแนวโน้มหรือประเมินสถานการณ์ต่างๆ เป็นสิ่งที่มีคุณค่าอย่างมาก อย่างไรก็ตาม ข้อมูลที่ไม่เกี่ยวข้องหรือซ้ำซ้อนอาจส่งผลเสียต่อการวิเคราะห์และลดความแม่นยำของผลลัพธ์ โดยเฉพาะในงานที่เกี่ยวข้องกับการวิเคราะห์ความคิดเห็นหรือความรู้สึกของผู้คน เช่น การวิเคราะห์ความรู้สึก (Sentiment Analysis) ซึ่งข้อมูลที่ไม่เกี่ยวข้องหรือผิดพลาดสามารถทำให้ผลการวิเคราะห์ไม่แม่นยำ ดังนั้น การคัดเลือกข้อมูลที่มีคุณภาพและเกี่ยวข้องจึงเป็นสิ่งสำคัญที่จะช่วยให้ได้ผลลัพธ์ที่น่าเชื่อถือและมีประสิทธิภาพ นอกจากนี้ยังมีผลต่อการประหยัดทรัพยากรในด้านการประมวลผลข้อมูลและการพัฒนาโมเดลที่มีความแม่นยำและประสิทธิภาพสูงสุด ซึ่งเป็นปัจจัยสำคัญในการพัฒนาเทคโนโลยีต่างๆ โดยเฉพาะในอุตสาหกรรมที่ใช้ข้อมูลในการตัดสินใจเชิงธุรกิจหรือบริการต่างๆ

การคัดเลือกข้อมูลและการทำป้ายกำกับข้อมูล (Data Labeling) เป็นกระบวนการที่มีต้นทุนสูง โดยเฉพาะเมื่อมีข้อมูลจำนวนมากหรือจำเป็นต้องใช้ผู้เชี่ยวชาญเฉพาะด้านเพื่อให้การทำป้ายกำกับข้อมูลถูกต้อง ซึ่งการคัดเลือกข้อมูลที่เหมาะสมสามารถช่วยลดต้นทุนและเวลาในการฝึกแบบจำลองได้ โดยลดเวลาหรือทรัพยากรในการทำป้ายกำกับข้อมูลกับข้อมูลที่ไม่จำเป็น ซึ่งการเรียนรู้เชิงรุก (Active Learning) เป็นหนึ่งในวิธีที่นิยมใช้เพื่อคัดเลือกข้อมูลที่มีความสำคัญที่สุดในการฝึกแบบจำลอง โดยช่วยให้สามารถเลือกข้อมูลที่คาดว่าจะมีประโยชน์มากที่สุดได้ในกระบวนการฝึกแบบจำลอง การเรียนรู้เชิงรุกไม่เพียงแต่ช่วยให้สามารถคัดเลือกข้อมูลที่มีความสำคัญสูงสุดในการฝึกแบบจำลอง แต่ยังช่วยเพิ่มความสามารถในการปรับปรุงผลลัพธ์ของโมเดลให้มีความแม่นยำมากยิ่งขึ้นโดยการคัดเลือกข้อมูลที่มีคุณภาพสูงมาใช้ในงานในกระบวนการเรียนรู้

งานวิจัยนี้ได้นำเสนอการใช้การเรียนรู้เชิงรุกในการคัดเลือกข้อมูลเพื่อใช้ในการวิเคราะห์ความรู้สึกจากข้อความที่เกี่ยวข้องกับการท่องเที่ยวในกรุงเทพมหานครจากทวิตเตอร์ โดยแบ่งความรู้สึกออกเป็นสามประเภท ได้แก่ เชิงบวก เป็นกลาง และเชิงลบ ข้อความที่ไม่เกี่ยวข้อง เช่น ข่าวสารหรือข้อความสแปม (Spam) จะไม่ได้รับการทำป้ายกำกับและจะถูกคัดออกจากกระบวนการ ทั้งนี้ งานวิจัยนี้ได้ใช้วิธีการคัดเลือกข้อมูลหลายรูปแบบ เช่น วิธีการเลือกข้อมูลแบบละโมภ (Greedy Approach) ซึ่งเน้นเลือกข้อมูลที่คาดว่าจะมีประโยชน์สูงสุด วิธีการสุ่มตัวอย่างแบบทอมป์สัน (Thompson Sampling) ที่ใช้การประมาณค่าความน่าจะเป็นในการเลือกข้อมูลที่มีประโยชน์ และการสุ่มตัวอย่างแบบสุ่ม (Random Sampling) เพื่อเปรียบเทียบประสิทธิภาพของแต่ละวิธีในการคัดเลือกข้อมูลที่เกี่ยวข้อง โดยการเลือกข้อมูลที่มีประโยชน์ต่อการวิเคราะห์และลดข้อมูลที่เป็นขยะจะช่วยให้เพิ่มความแม่นยำและประสิทธิภาพของแบบจำลองในการวิเคราะห์ความรู้สึก การเปรียบเทียบเหล่านี้ยังสามารถเปิดมุมมองใหม่ในการพัฒนาวิธีการคัดเลือกข้อมูลในอนาคตเพื่อเพิ่มประสิทธิภาพในการนำข้อมูลไปใช้ในการประมวลผลและวิเคราะห์ในบริบทที่หลากหลายมากยิ่งขึ้น

การประยุกต์ใช้การเรียนรู้เชิงรุกในงานวิจัยนี้มีเป้าหมายเพื่อแสดงให้เห็นถึงประโยชน์ของการเลือกข้อมูลที่มีคุณภาพสูงในการฝึกแบบจำลองการวิเคราะห์ความรู้สึก ข้อมูลที่ได้รับการคัดเลือกและติดป้ายกำกับอย่างเหมาะสมจะช่วยให้ตัวแบบสามารถเรียนรู้ได้ดีขึ้นและเพิ่มความน่าเชื่อถือของผลลัพธ์ที่ได้ การปรับปรุงกระบวนการคัดเลือกข้อมูลให้มีความแม่นยำยิ่งขึ้นไม่เพียงแต่จะช่วยในงานวิเคราะห์ข้อมูลข้อความเท่านั้น แต่ยังสามารถนำไปปรับใช้กับข้อมูลประเภทอื่นๆ เช่น ข้อมูลภาพหรือวิดีโอในอนาคตได้ด้วย โดยเฉพาะในงานที่ต้องการความแม่นยำในการวิเคราะห์และตัดสินใจจากข้อมูลจำนวนมาก การพัฒนาและปรับใช้วิธีการคัดเลือกข้อมูลที่เหมาะสมจึงถือเป็นหนึ่งในกระบวนการสำคัญที่สามารถยกระดับประสิทธิภาพการทำงานในหลายๆ ด้านได้

การทบทวนวรรณกรรม

การวิจัยในครั้งนี้ครอบคลุมทฤษฎีสำคัญ 5 ประการ ได้แก่ Term Frequency-Inverse Document Frequency (TF-IDF), การวิเคราะห์ถดถอยโลจิสติก (Logistic Regression), วิธีการเลือกข้อมูลโดยวิธีละโมภ (Greedy Approach),

Thompson Sampling (TS), และการประมาณแบบลาปลาซ (Laplace Approximation; LA) ซึ่งเกี่ยวข้องโดยตรงกับปัญหาการเรียนรู้เชิงรุกในข้อมูลประเภทข้อความ โดย TF-IDF ใช้ในการประมวลผลข้อความเพื่อประเมินความสำคัญของคำในเอกสาร, การวิเคราะห์ถดถอยโลจิสติก ใช้คำนวณความน่าจะเป็นของผลลัพธ์ที่เป็นไปได้จากคุณลักษณะของข้อมูล, Greedy Approach เลือกข้อมูลที่ให้ผลตอบแทนสูงสุดในแต่ละรอบโดยไม่คำนึงถึงผลลัพธ์ในอนาคต, และ Thompson Sampling (TS) ใช้ในการตัดสินใจเพื่อแสวงหาผลลัพธ์ที่ดีที่สุดโดยการเลือกการกระทำที่ให้ข้อมูลมากที่สุด โดยใช้ทฤษฎีความน่าจะเป็นในการคำนวณความไม่แน่นอนจากข้อมูลที่ผ่านมา

TF-IDF เป็นเทคนิคที่ใช้กันอย่างแพร่หลายในการดึงข้อมูลจากเอกสาร (Information retrieval) และการทำเหมืองข้อความ (Text mining) ซึ่งถูกนิยามใน Aizawa (2003) ใช้เพื่อกำหนดความสำคัญของคำในเอกสารด้วยการวัดทางสถิติที่ประเมินความเกี่ยวข้องของคำกับเอกสารในชุดเอกสาร โดยคำนวณจากความถี่ของแต่ละคำศัพท์ในเอกสารและเปรียบเทียบกับความถี่ของคำศัพท์เดียวกันในเอกสารทั้งหมด วิธีการนี้กำหนดน้ำหนักที่สูงให้กับคำที่ใช้อยู่ในเอกสารนั้นแต่พบน้อยในเอกสารอื่นๆ ในทางกลับกันจะกำหนดน้ำหนักที่ต่ำให้กับคำที่ใช้กันทั่วไปในเอกสารทั้งหมด TF-IDF เคยถูกใช้ในเครื่องมือค้นหา (Search engine) ใน Xu (2014) การจัดกลุ่มเอกสาร (Document clustering) ใน Bafna, Pramod and Vaidya (2016) และการจัดประเภทเอกสาร (Document classification) ใน Hakim et al. (2014) วิธีการคำนวณความสำคัญของคำศัพท์ด้วยวิธี TF-IDF เกิดจากการคำนวณ Term Frequency (TF) คูณด้วย Inverse Document Frequency (IDF) ซึ่งการคำนวณ TF ประเมินจากความเกี่ยวข้องของคำต่อเอกสารนั้นๆ โดยค่าของ Term Frequency เป็นค่าที่บอกความถี่ของคำแต่ละคำที่ปรากฏในเอกสารหนึ่งๆ โดยคิดคำนวณจากการนำจำนวนครั้งที่คำนั้นๆ ปรากฏในเอกสารมาหารด้วยจำนวนคำทั้งหมดในเอกสารนั้นๆ ในขณะที่ Inverse Document Frequency (IDF) เป็นการคำนวณค่าน้ำหนักความสำคัญของแต่ละคำในเอกสารทั้งหมด ค่า IDF ที่ต่ำบ่งบอกว่าคำเหล่านั้นไม่สามารถดึงจุดเด่นของเอกสารที่คำเหล่านั้นปรากฏอยู่ออกมาได้ ค่า IDF สามารถคำนวณได้จากสมการใน Aizawa (2003) การวิเคราะห์ถดถอยโลจิสติกเป็นตัวแทนโมเดลที่ใช้สำหรับการจำแนกประเภท (Classification) ในกรณีที่ผลลัพธ์ที่ทำนายเป็นค่าความน่าจะเป็น (Probability) ของแต่ละคลาสที่เป็นไปได้ โดยปกติแล้วการวิเคราะห์ถดถอยโลจิสติกจะใช้ในการทำนายเหตุการณ์ที่มีสองผลลัพธ์ (Binary Outcomes) ตาม Hosmer, Lemeshow and Sturdivant (2000) ซึ่งสำหรับงานวิจัยนี้เรามุ่งเน้นไปที่การทำนายข้อมูลที่มีประโยชน์ (Non-garbage) และข้อมูลที่เป็นขยะ (Garbage) การวิเคราะห์ถดถอยโลจิสติกจะทำการเรียนรู้ฟังก์ชันที่สามารถคำนวณความน่าจะเป็นที่เหตุการณ์ใดเหตุการณ์หนึ่งจะเกิดขึ้น โดยใช้สมการของ Logistic Function หรือ Sigmoid Function

วิธีการเลือกข้อมูลแบบละโมภ (Greedy Approach) เป็นกระบวนการที่เลือกการกระทำหรือการเลือกข้อมูลที่ทำให้ค่าผลตอบแทนสูงสุดในแต่ละขั้นตอนการทำงาน โดยไม่พิจารณาผลลัพธ์ในอนาคต การเลือกการกระทำในแต่ละรอบจะเลือกค่าผลตอบแทนสูงสุดจากตัวเลือกที่มีอยู่ในขณะนั้น โดยมุ่งเน้นการหาผลตอบแทนที่ดีที่สุดในระยะสั้น โดยไม่คำนึงถึงการสำรวจการกระทำที่อาจจะให้ผลในระยะยาว แม้ว่าจะไม่มีการพิจารณาความไม่แน่นอนหรือการสำรวจในระยะยาว แต่การเลือกแบบละโมภสามารถทำให้ได้ผลลัพธ์ที่รวดเร็วในกรณีที่การกระทำต่างๆ มีผลตอบแทนที่ชัดเจนและไม่เปลี่ยนแปลงมากนัก ซึ่งทำให้ได้ผลลัพธ์ที่ดีในระยะสั้นได้อย่างรวดเร็ว เช่น Koller and Friedman (2009)

Thompson Sampling (TS) เป็นวิธีการตัดสินใจในการเลือกการกระทำเพื่อแสวงหาผลประโยชน์ภายใต้การแจกแจงของสิ่งแวดล้อมของปัญหา โดยเลือกการกระทำที่ให้ค่าคาดหวังสูงที่สุดตามตัวอย่างที่สุ่มเลือกมาเพื่อนำมาการประมาณหาการแจกแจงหลังของสิ่งแวดล้อมที่เหมาะสมที่สุด

การประมาณแบบลาปลาซ (Laplace Approximation; LA) เป็นวิธีที่ใช้ในการประมาณการแจกแจงหลังในรูปแบบเกาส์เซียน โดยค่าเฉลี่ยของการแจกแจงหลังได้จากการคำนวณ Bayesian logistic regression ผ่านข้อมูลที่ถูกเลือก ซึ่งเคยถูกพิสูจน์ในบทที่ 6.2.3 ของ Bishop and Nasrabadi (2006) และความแปรปรวนของการกระจายหลังคือส่วนผกผันของเมทริกซ์เฮสเซียน (Hessian matrix) ของลอการิธึมของการแจกแจงหลังที่ได้จากการค้นหาค่าประมาณค่าหลังสุด (Bishop & Nasrabadi, 2006)

วิธีดำเนินการวิจัย

เงื่อนไขการแบ่งข้อมูล

ในงานวิจัยนี้ที่มุ่งเปรียบเทียบประสิทธิภาพการคัดเลือกข้อมูลที่มีประโยชน์ (Non-garbage) เพื่อทำป้ายกำกับในแต่ละวิธี ซึ่งขอบเขตการคัดเลือกข้อมูล คือการคัดเลือกทวิตเกี่ยวกับการท่องเที่ยวของจังหวัดกรุงเทพมหานครในช่วงกรกฎาคม-ธันวาคม 2563 เพื่อมาทำการวิเคราะห์ความรู้สึก โดยแบ่งทวิตที่มีประโยชน์ต่อการวิเคราะห์คือทวิตที่แสดงความรู้สึกออกมาประเภทเดียว และทวิตที่เป็นขยะคือทวิตที่แสดงความรู้สึกมากกว่าหนึ่งความรู้สึกหรือทวิตที่เป็นข่าวและสแปม ซึ่งการแบ่งข้อมูลในช่วงเริ่มต้นเพื่อให้ตัวแบบการวิเคราะห์ถดถอยโลจิสติกเรียนรู้เป็นขั้นตอนสำคัญที่มีผลต่อการคัดเลือกข้อมูลในขั้นถัดไปที่ตัวแบบจะเลือกเข้ามาทำป้ายกำกับ ดังนั้นจึงได้เลือกวิธีการสุ่มในการแบ่งข้อมูลเริ่มต้นสำหรับฝึกตัวแบบ โดยแบ่งออกเป็นดังนี้

- 1) ทวิตในเชิงบวกเกี่ยวกับการท่องเที่ยวของกรุงเทพ 250 ทวิตจาก 4,895 ทวิต
- 2) ทวิตที่เป็นกลางเกี่ยวกับการท่องเที่ยวของกรุงเทพ 250 ทวิตจาก 7,170 ทวิต
- 3) ทวิตในเชิงลบเกี่ยวกับการท่องเที่ยวของกรุงเทพ 250 ทวิตจาก 1,685 ทวิต
- 4) ทวิตที่ไม่สามารถระบุความรู้สึกได้จำนวน 250 ทวิตจาก 48,095 ทวิต

ทวิตที่ไม่ได้ถูกเลือกในขั้นตอนนี้จะถูกเก็บไว้ในพูลซึ่งจะนำมาใช้ในการเลือกเพื่อทำป้ายกำกับในขั้นตอนถัดไป โดยมีทวิตที่มีประโยชน์ทั้งหมด 13,000 ทวิต และทวิตที่เป็นขยะทั้งหมด 47,845 ทวิต

เงื่อนไขการฝึกฝนตัวแบบ

การฝึกฝนตัวแบบโมเดลการวิเคราะห์ถดถอยโลจิสติกเริ่มต้นจากการเรียนรู้ผ่านทวิต 1,000 ทวิตที่ถูกแบ่งไว้ตามเงื่อนไขการแบ่งข้อมูล จากนั้นตัวแบบจะทำการเลือกข้อมูลจากพูลมาครั้งละ 250 ทวิต โดยใช้วิธีการเลือกข้อมูลทั้งสามวิธี ดังนี้

- 1) วิธีการเลือกข้อมูลแบบสุ่ม (Random Approach) เป็นการเลือกทวิตจากพูลมาแบบสุ่มเพื่อให้โมเดลเรียนรู้ โดยไม่พิจารณาถึงความน่าจะเป็นหรือคุณลักษณะใดๆ ของทวิต
- 2) วิธีการเลือกข้อมูลโดยวิธีละโมภ (Greedy Approach) เป็นการเลือกทวิตจากพูลมาโดยเลือกทวิตที่ตัวแบบคำนวณความน่าจะเป็นที่จะไม่เป็นขยะ (Non-garbage) สูงที่สุด

Algorithm 1 Greedy Approach

```

1: for t = 0, 1, 2, 3, ... do
2:    $R_{t+1,a} \leftarrow$  Logistic Regression Probability
3:    $A_t \leftarrow \operatorname{argmax}_{a \in A} R_{t+1,a}$            # select action
4:
5:   Apply  $A_t$  and observe  $O_{t+1}$ 
6:    $H_{t+1} \leftarrow \operatorname{append}(H_t, A_t, O_{t+1}))$    # update history
7:
8:   Logistic Regression trained with  $H_{t+1}$ 
9:
10: end for

```

- 3) วิธีการเลือกข้อมูลโดยการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาปลาซ (Thompson Sampling using Laplace Approximation) เป็นการเลือกทวิตจากพูลที่ตัวแบบคำนวณว่ามีความน่าจะเป็นที่ทวิตนั้นจะไม่เป็นขยะสูง

ที่สุดภายใต้การสุ่มตัวอย่างของทอมป์สัน โดยการสุ่มค่าสัมประสิทธิ์จาก μ และ Σ ภายใต้การแจกแจงแบบปกติ (Normal Distribution)

Algorithm 2 Thompson Sampling Algorithm

```

1: for  $t = 0, 1, 2, 3, \dots$  do
2:    $\mu \leftarrow$  Logistic Regression Coefficient
3:    $\Sigma \leftarrow$  diagonal Inverse Hessian
4:   Sample  $\hat{\theta}_t \sim P_t$  # sample model
5:
6:    $A_t \leftarrow \operatorname{argmax}_{a \in A} R_{t+1,a}$  # select action
7:   Apply  $A_t$  and observe  $O_{t+1}$ 
8:    $H_{t+1} \leftarrow \operatorname{append}(H_t, A_t, O_{t+1})$  # update history
9:
10:   $P_{t+1} \leftarrow \mathbb{P}(\theta | H_{t+1})$  # update distribution
11: end for
  
```

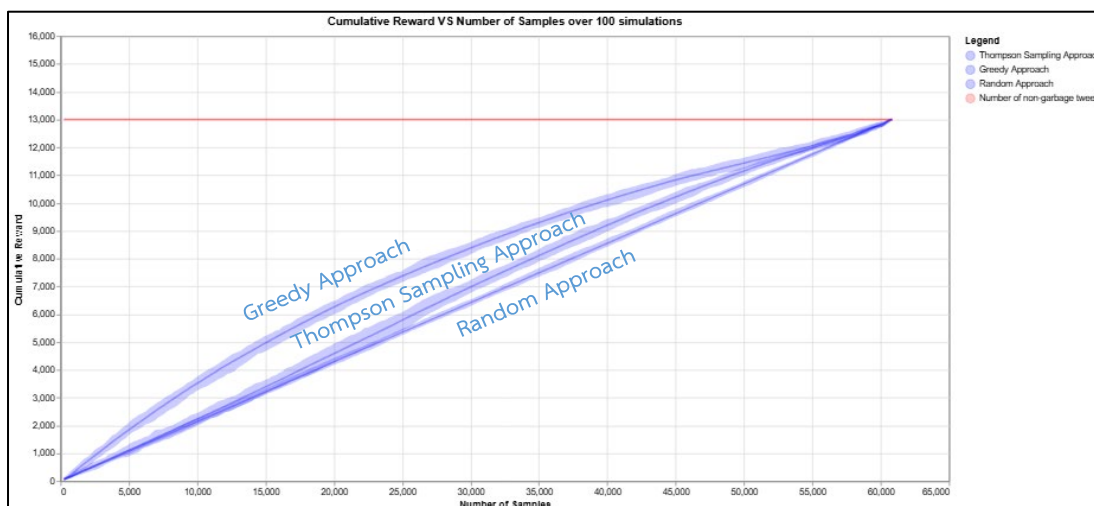
ขั้นตอนการวิจัย

การดำเนินงานวิจัยเริ่มต้นจากการเตรียมข้อมูลและการทำป้ายกำกับข้อมูลทวิตตามความรู้สึกของผู้อ่าน ต่อมาข้อมูลจะได้รับการทำความสะอาดเพื่อลบเครื่องหมายพิเศษ ลิงค์ และคำที่เป็น stop word หลังจากนั้นข้อมูลจะถูกแปลงเป็นตัวเลขโดยใช้ TF-IDF vectorizer และแบ่งข้อมูลเพื่อฝึกตัวแบบด้วยข้อมูลเริ่มต้น 1,000 ทวิต โดยมีข้อมูลในพูลรวมทั้งสิ้น 60,845 ทวิต หลังจากนั้นตัวแบบการวิเคราะห์ถดถอยโลจิสติกจะได้รับการฝึกด้วยข้อมูลเริ่มต้นทั้ง 1,000 ทวิตที่เตรียมไว้ และทำการคัดเลือกข้อมูลด้วยสามวิธีในการคัดเลือกข้อมูล ดังนี้

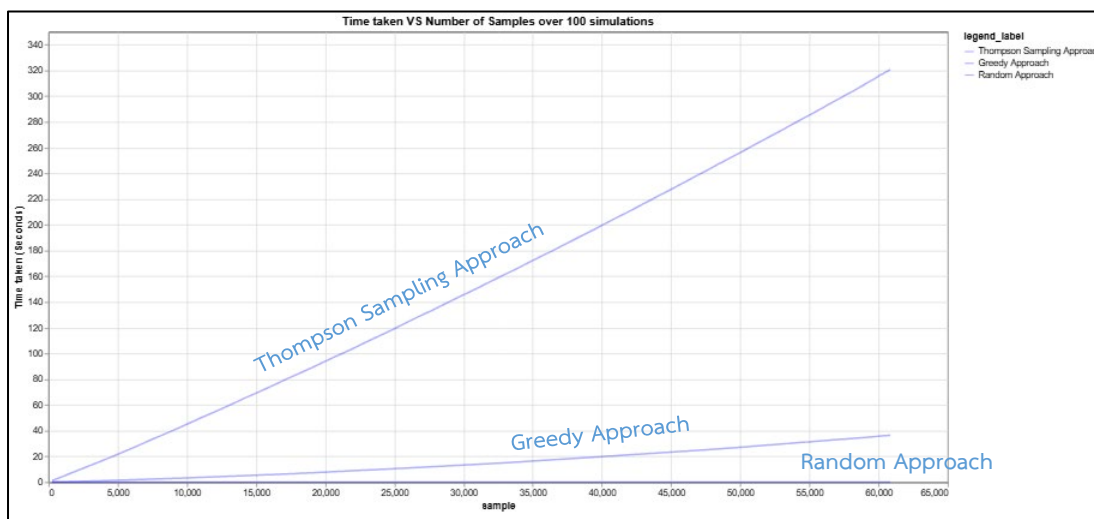
- 1) การเลือกข้อมูลแบบสุ่ม เลือกข้อมูลจากในพูลแบบสุ่มโดยไม่พิจารณาความน่าจะเป็นหรือคุณลักษณะใดๆ ของทวิต
- 2) การเลือกข้อมูลโดยวิธีละโมภ เลือกข้อมูลโดยตัวแบบการวิเคราะห์ถดถอยโลจิสติกจะคำนวณความน่าจะเป็นที่ข้อมูลไม่ใช่ขยะสำหรับทวิตแต่ละทวิตในพูล จากนั้นตัวแบบจะเลือกทวิตที่มีความน่าจะเป็นสูงที่สุด
- 3) การเลือกข้อมูลโดยการสุ่มตัวอย่างของทอมป์สัน เลือกข้อมูลโดยตัวแบบการวิเคราะห์ถดถอยโลจิสติกจะถูกประมาณการแจกแจงหลังด้วยการประมาณแบบลาพลาซ จากนั้นตัวแบบการวิเคราะห์ถดถอยโลจิสติกจะถูกสุ่มค่าสัมประสิทธิ์ภายใต้การแจกแจงหลังในแต่ละรอบ ตัวแบบการวิเคราะห์ถดถอยโลจิสติกใหม่จะคำนวณความน่าจะเป็นที่ข้อมูลไม่ใช่ขยะสำหรับทวิตแต่ละทวิตในพูล จากนั้นตัวแบบจะเลือกทวิตที่มีความน่าจะเป็นสูงที่สุด

โดยแต่ละวิธีจะทำการเลือกทวิตจากพูล 250 ทวิตในแต่ละรอบการเรียนรู้ และตัวแบบการวิเคราะห์ถดถอยโลจิสติกจะถูกปรับด้วยข้อมูลใหม่ที่คัดเลือกในแต่ละรอบ จนข้อมูลในพูลหมด การทดลองจะถูกทำซ้ำ 100 รอบด้วยข้อมูลเริ่มต้น 1,000 ทวิตที่แตกต่างกันเพื่อเก็บผลการทดลอง โดยมีการเปรียบเทียบประสิทธิภาพของโมเดลทั้งในแง่ของจำนวนข้อมูลที่มีประโยชน์ที่แต่ละวิธีเลือกมา และเวลาในการฝึกโมเดล ทั้งสามวิธีจะได้รับการเปรียบเทียบเพื่อสรุปผลการทดลองและอภิปรายผลในที่สุด

ผลการวิจัย



ภาพที่ 1 ผลการทดลองเฉลี่ยทั้ง 100 รอบในการเลือกทวิตเพื่อมาทำป้ายกำกับและเรียนรู้ผ่านตัวแบบการวิเคราะห์ถดถอยโลจิสติกด้วยวิธีการเลือกข้อมูลแบบสุ่ม, วิธีการเลือกข้อมูลโดยวิธีละโมบ, และวิธีการสุ่มตัวอย่างของทอมป์สัน ด้วยการประมาณค่าแบบลาพลาซ แกน x คือ จำนวนทวิตที่ถูกเลือกทั้งหมด แกน y คือ จำนวนทวิตที่ไม่ใช่ขยะที่ถูกเลือก



ภาพที่ 2 ผลการทดลองเฉลี่ยทั้ง 100 รอบในการเลือกทวิตเพื่อมาทำป้ายกำกับและเรียนรู้ผ่านตัวแบบการวิเคราะห์ถดถอยโลจิสติกด้วยวิธีการเลือกข้อมูลแบบสุ่ม, วิธีการเลือกข้อมูลโดยวิธีละโมบ, และวิธีการสุ่มตัวอย่างของทอมป์สัน ด้วยการประมาณค่าแบบลาพลาซ แกน x คือ จำนวนทวิตที่ถูกเลือกทั้งหมด แกน y คือ เวลาที่ตัวแบบใช้ในการเลือกและเรียนรู้ข้อมูล

ผลการวิจัยด้วยวิธีการเลือกข้อมูลแบบสุ่ม (Random Approach) พบว่า ในการทดลอง 100 รอบ ตัวแบบไม่มีการเรียนรู้จากทวิตที่ถูกเลือกมา ส่งผลให้ประสิทธิภาพในการเลือกทวิตที่ไม่ใช่ขยะคงที่ตลอดการทดลอง โดยจำนวนทวิตที่ไม่ใช่ขยะที่ถูกเลือกในแต่ละรอบอยู่ที่ประมาณ 53 ทวิต ในขณะที่เวลาเฉลี่ยที่ใช้ในการเลือกทวิต 250 ทวิตอยู่ที่ประมาณ 0.316 มิลลิวินาทีต่อรอบ ซึ่งค่อนข้างคงที่ตลอดการทดลอง

ในส่วนของวิธีการเลือกข้อมูลโดยวิธีละโมบ (Greedy Approach) การทดลองแสดงให้เห็นว่าในช่วงครึ่งแรกของการเลือกทวิต ประสิทธิภาพในการเลือกทวิตที่ไม่ใช่ขยะดีขึ้น โดยเฉลี่ยจะอยู่ที่ประมาณ 69 ทวิตต่อรอบ แต่จะลดลงเหลือ

ประมาณ 37 ทวีตในช่วงครึ่งหลังของการทดลอง เนื่องจากจำนวนทวีตในพูลลดลง ในขณะที่เดียวกันเวลาในการเลือกทวีตและเรียนรู้ข้อมูลมีแนวโน้มเพิ่มขึ้นโดยเฉลี่ยจะอยู่ที่ประมาณ 0.151 วินาทีต่อรอบ

สำหรับวิธีการเลือกข้อมูลโดยการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซ (Thompson Sampling using Laplace Approximation) ผลการทดลองพบว่า ประสิทธิภาพในการเลือกทวีตที่ไม่ใช่ขยะจะคงที่จนถึงช่วงท้ายของการทดลอง โดยในช่วงครึ่งแรกจะเลือกทวีตที่ไม่ใช่ขยะเฉลี่ยประมาณ 58 ทวีต และลดลงเหลือประมาณ 48 ทวีต ในช่วงครึ่งหลัง ตัวแบบใช้เวลาในการเรียนรู้นานขึ้นเมื่อเทียบกับวิธีการเลือกข้อมูลแบบอื่นๆ โดยเวลาเฉลี่ยที่ใช้ในการเลือกทวีต 250 ทวีตอยู่ที่ประมาณ 1.320 วินาทีต่อรอบ

สรุปและอภิปรายผลการวิจัย

ประสิทธิภาพการคัดเลือกข้อมูลที่เป็นประโยชน์ด้วยวิธีการเลือกข้อมูลแบบสุ่ม การเลือกข้อมูลโดยวิธีละโมบ และการเลือกข้อมูลโดยการสุ่มตัวอย่างของทอมป์สันสามารถอภิปรายผลได้ ดังนี้

1) การเปรียบเทียบประสิทธิภาพการเลือกข้อมูลทั้งสามวิธีแสดงให้เห็นว่าวิธีการเลือกข้อมูลโดยวิธีละโมบ มีประสิทธิภาพในการเลือกทวีตที่ไม่ใช่ขยะสูงที่สุดเมื่อเทียบกับวิธีอื่นๆ เนื่องจากมีการเรียนรู้จากทวีตที่เลือกมาด้วยการวิเคราะห์แบบถดถอยโลจิสติกและปรับปรุงตัวแบบอย่างต่อเนื่อง ทำให้วิธีนี้มีประสิทธิภาพดีที่สุดในทุกช่วงของการทดลอง ในทางตรงกันข้ามวิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซ มีการปรับปรุงประสิทธิภาพหลังจากการทดลองไปได้สักระยะ โดยในช่วงแรกนั้นประสิทธิภาพคล้ายกับการเลือกข้อมูลแบบสุ่ม เนื่องจากตัวแบบเริ่มเรียนรู้และสร้างการแจกแจงหลังของทวีตเมื่อมีข้อมูลเพียงพอแล้ว ดังนั้นในช่วงท้ายๆ ประสิทธิภาพของวิธีนี้จะดีขึ้นอย่างเห็นได้ชัด แต่ยังคงต่ำกว่าวิธีการเลือกข้อมูลโดยวิธีละโมบ ส่วนวิธีการเลือกข้อมูลแบบสุ่ม ไม่ได้ทำการเรียนรู้จากทวีตที่เลือกมา ทำให้ประสิทธิภาพคงที่ตลอดทั้งการทดลอง และมีประสิทธิภาพต่ำที่สุดเมื่อเทียบกับวิธีอื่นๆ

2) การเปรียบเทียบเวลาที่ใช้ในการเลือกข้อมูลทั้งสามวิธี เมื่อพิจารณาในแง่ของเวลา วิธีการเลือกข้อมูลแบบสุ่ม ใช้เวลาน้อยที่สุด เนื่องจากไม่ต้องทำการเรียนรู้หรือปรับปรุงตัวแบบจากทวีตที่เลือกมา ทำให้มีประสิทธิภาพในการเลือกข้อมูลที่รวดเร็วที่สุด ในขณะที่ วิธีการเลือกข้อมูลโดยวิธีละโมบ ใช้เวลามากขึ้น เนื่องจากตัวแบบต้องทำการเรียนรู้จากทวีตที่เลือกมาและปรับปรุงโมเดลอย่างต่อเนื่อง ซึ่งทำให้เวลาที่ใช้ในการเลือกและเรียนรู้ข้อมูลเพิ่มขึ้นตามจำนวนข้อมูลที่มี สำหรับ วิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซ ใช้เวลานานที่สุด เนื่องจากต้องทำการเรียนรู้และปรับปรุงโมเดลจากทวีตที่เลือกมา รวมถึงต้องใช้เวลาในการประมาณค่าเฉลี่ยและแปรปรวนของการแจกแจงหลัง ซึ่งจะทำให้วิธีนี้ใช้เวลามากกว่าวิธีอื่นๆ โดยเฉพาะในขั้นตอนที่ต้องประมาณค่าทางสถิติจากข้อมูลที่มี

3) การสรุปความแตกต่างและประสิทธิภาพของทั้งสามวิธีจากการทดลองทั้ง 100 รอบในการเลือกทวีตเพื่อมาทำป้ายกำกับและเรียนรู้ผ่านตัวแบบการวิเคราะห์ถดถอยโลจิสติก พบว่าตัวแบบวิธีการเลือกข้อมูลแบบสุ่มใช้เวลาในการเลือกเฉลี่ยประมาณ 0.316 มิลลิวินาทีต่อรอบ และประสิทธิภาพของตัวแบบไม่ดีขึ้นเนื่องจากไม่มีการเรียนรู้หรือปรับปรุงจากทวีตที่เลือกมา ตัวแบบวิธีการเลือกข้อมูลโดยวิธีละโมบใช้เวลาในการเลือกเฉลี่ยประมาณ 0.151 วินาทีต่อรอบ ซึ่งสูงกว่าวิธีเลือกแบบสุ่มเนื่องจากต้องเรียนรู้และปรับปรุงโมเดลจากทวีตที่เลือกมาโดยตรง ประสิทธิภาพของตัวแบบสูงขึ้นอย่างเห็นได้ชัดในช่วงแรก แต่ในช่วงหลังประสิทธิภาพของตัวแบบจะลดลงเนื่องจากทวีตที่ไม่ใช่ขยะในพูลมีจำนวนน้อยลง ตัวแบบวิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซใช้เวลาในการเลือกเฉลี่ยประมาณ 1.320 วินาทีต่อรอบ ซึ่งสูงกว่าวิธีเลือกแบบสุ่มและวิธีการเลือกข้อมูลโดยวิธีละโมบเนื่องจากต้องใช้เวลาในการเรียนรู้และปรับปรุงโมเดลจากทวีตที่เลือกมา นอกจากนี้ยังต้องใช้เวลาในการประมาณค่าเฉลี่ยและแปรปรวนของการแจกแจงหลัง ซึ่งวิธีนี้มีประสิทธิภาพที่ต่ำในช่วงแรกเนื่องจากการเรียนรู้หรือปรับปรุงตัวแบบจากทวีตที่เลือกมาเป็นการประมาณการแจกแจงหลังของทวีตที่เรียนรู้ แต่หลังจากการทดลองผ่านไป ตัวแบบเริ่มแม่นยำขึ้นและมีประสิทธิภาพเพิ่มขึ้นอย่างชัดเจน

โดยสรุปแล้ว วิธีการเลือกข้อมูลโดยวิธีละโมบ มีประสิทธิภาพสูงที่สุดตลอดรอบการทดลอง ส่วนวิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซ จะมีประสิทธิภาพใกล้เคียงกับ วิธีการเลือกข้อมูลแบบสุ่ม ในช่วงแรก แต่จะดีขึ้นในช่วงหลังของการทดลอง

ข้อเสนอแนะในการวิจัยครั้งต่อไป

1) การประยุกต์ใช้วิธีการเลือกข้อมูลแบบละโมบและวิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซ ในกระบวนการสุ่มตัวอย่างแบบพูล เพื่อคัดเลือกข้อมูลที่มีประโยชน์สูงสุดสำหรับการทำป้ายกำกับ การใช้ตัวแบบการเลือกข้อมูลทั้งสองวิธีเป็นแนวทางที่เหมาะสมสำหรับการคัดเลือกข้อมูลประเภทข้อความที่มีประโยชน์สูงสุดสำหรับการทำป้ายกำกับทำให้ลดค่าใช้จ่ายในการทำป้ายกำกับ สำหรับงานวิจัยในบริบทอื่นที่ต้องใช้ความเชี่ยวชาญหรือจำนวนข้อมูลในการทำป้ายกำกับปริมาณมาก เช่น ข้อมูลประเภทวิดีโอ หรือเสียง เป็นต้น สองวิธีนี้เป็นอีกทางเลือกสำหรับการคัดเลือกข้อมูล

2) การวิจัยเพิ่มเติมเกี่ยวกับตัวแบบวิธีการเลือกข้อมูลแบบละโมบและวิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซ

การวิจัยนี้เน้นไปที่การเปรียบเทียบระหว่างตัวแบบวิธีการเลือกข้อมูลแบบสุ่ม, วิธีการเลือกข้อมูลแบบละโมบ, และวิธีการสุ่มตัวอย่างของทอมป์สันด้วยการประมาณค่าแบบลาพลาซด้วยพหุขนาดจำกัด แต่ในโลกความเป็นจริงข้อมูลมีจำนวนมากมหาศาล ทำให้การเปรียบเทียบและพัฒนาตัวแบบวิธีการเลือกข้อมูลด้วยจำนวนที่มากขึ้น ทำให้เห็นผลในระยะยาวที่แม่นยำได้ การประยุกต์ใช้วิธีการเลือกข้อมูลทั้งสองวิธีกับงานวิจัยในบริบทหรือข้อมูลประเภทต่างๆ สามารถแสดงให้เห็นประสิทธิภาพที่แตกต่างจากในงานวิจัยนี้ ขึ้นอยู่กับความซับซ้อนและจำนวนของข้อมูล

เอกสารอ้างอิง

- Aizawa, A. (2003). An information-theoretic perspective of tf-idf measures. *Information Processing & Management*, 39(1), 45-65.
- Bafna, P., Pramod, D., & Vaidya, A. (2016). *Document clustering: TF-IDF approach*. 2016 International Conference on Electrical, Electronics, and Optimization Techniques (ICEEOT),
- Bishop, C. M., & Nasrabadi, N. M. (2006). *Pattern recognition and machine learning* (Vol. 4). Springer.
- Hakim, A. A., Erwin, A., Eng, K. I., Galinium, M., & Muliady, W. (2014). *Automated document classification for news article in Bahasa Indonesia based on term frequency inverse document frequency (TF-IDF) approach*. 2014 6th international conference on information technology and electrical engineering (ICITEE),
- Hosmer, D. W., Lemeshow, S., & Sturdivant, R. X. (2000). *Applied logistic regression*. Hoboken. In: NJ: John Wiley & Sons.
- Koller, D., & Friedman, N. (2009). *Probabilistic graphical models: principles and techniques*. MIT press.
- Murphy, K. P. (2012). *Machine learning: a probabilistic perspective*. MIT press.
- Russo, D. J., Van Roy, B., Kazerouni, A., Osband, I., & Wen, Z. (2018). A tutorial on thompson sampling. *Foundations and Trends® in Machine Learning*, 11(1), 1-96.
- Settles, B. (2009). *Active learning literature survey*.
- Xu, R. (2014). *POS weighted TF-IDF algorithm and its application for an MOOC search engine*. 2014 International Conference on Audio, Language and Image Processing.

Data Availability Statement: The raw data supporting the conclusions of this article will be made available by the authors, without undue reservation.

Conflicts of Interest: The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.



Copyright: © 2025 by the authors. This is a fully open-access article distributed under the terms of the Attribution-NonCommercial-NoDerivatives 4.0 International (CC BY-NC-ND 4.0).