

การพัฒนาประสิทธิภาพตัวแบบค้นหาเชิงความหมายจากเทคนิค Doc2Vec

เพื่อประยุกต์ใช้สำหรับการสืบค้นเพลงด้วยเสียงร้อง

Development of Semantic Search Model Based on The Doc2Vec Technique for an Application for Vocal Music Search

พิศาล สุขชี

Phisan Sookkhee

สาขาวิชาวิทยาการคอมพิวเตอร์ คณะศิลปศาสตร์และวิทยาศาสตร์ มหาวิทยาลัยราชภัฏศรีสะเกษ

Major of Computer Science, Faculty of Liberal Arts, Sisaket Rajabhat University

E-Mail: phisan.s@sskru.ac.th

บทคัดย่อ

การวิจัยครั้งนี้มีวัตถุประสงค์เพื่อ 1) พัฒนาประสิทธิภาพตัวแบบค้นหาเชิงความหมายจากเทคนิค Doc2Vec 2) พัฒนาระบบการค้นหาเพลงจากเสียงร้องด้วยตัวแบบการค้นหาเชิงความหมายจากเทคนิค Doc2Vec สำหรับข้อมูลที่ใช้ในการศึกษาวิจัยคือ เนื้อเพลงภาษาไทยประเภทเพลงลูกทุ่งจำนวน 1,500 เพลง ในการพัฒนาประสิทธิภาพของตัวแบบ ผู้วิจัยได้ทำการค้นหาคำไฮเปอร์พารามิเตอร์ที่เหมาะสมเพื่อเพิ่มประสิทธิภาพของตัวแบบ การค้นหาขนาดของคำสอบถามที่มีความเหมาะสม และส่งผลกระทบต่อความแม่นยำในการค้นหาของตัวแบบ ตลอดจนการนำตัวแบบที่ผ่านการพัฒนาประสิทธิภาพไปพัฒนาระบบต้นแบบเพื่อใช้สำหรับสืบค้นเพลงด้วยเสียงร้อง

ผลการวิจัยพบว่า 1) ผลการพัฒนาประสิทธิภาพตัวแบบค้นหาเชิงความหมายจากเทคนิค Doc2Vec พบว่า ทำให้ทราบขนาดของคำสอบถามที่ส่งผลกระทบต่อประสิทธิภาพการค้นหาของตัวแบบอยู่ที่ขนาดความยาวมากกว่า ร้อยละ 30 ของเนื้อเพลงต้นฉบับที่ต้องการสืบค้น และขนาดของเวกเตอร์ หรือมิติข้อมูลที่มีความเหมาะสมต่อการนำมาสร้างตัวแบบจากชุดข้อมูลเนื้อเพลงลูกทุ่งอยู่ที่ขนาดเท่ากับ 200 และการคงค่าหยุดไว้เนื้อเพลงจะส่งผลกระทบต่อประสิทธิภาพในการสืบค้นของตัวแบบสูงกว่าการกำกวมค่าหยุดออกไป 2) ผลการพัฒนาระบบการค้นหาเพลงจากเสียงร้องด้วยตัวแบบการค้นหาเชิงความหมายจากเทคนิค Doc2Vec พบว่าผู้สืบค้นสามารถค้นหาเพลงด้วยการป้อนเสียงร้องผ่านไมโครโฟน โดยเสียงร้องจะถูกแปลงให้เป็นข้อความด้วย Web Speech API นำไปผ่านกระบวนการตัดคำภาษาไทย และสร้างคำสอบถามเพื่อสืบค้นกับตัวแบบที่ถูกพัฒนาประสิทธิภาพให้เหมาะสม โดยจากการทดสอบซึ่งทำให้ผู้ใช้ทราบถึงความเหมาะสมว่าต้องร้องเนื้อเพลงด้วยความยาวเท่าใดที่จะทำให้สามารถสืบค้นบทเพลงตามที่ต้องการได้ โดยคำสอบถามนั้นไม่จำเป็นต้องมีความถูกต้องตามเนื้อหาของเพลงต้นฉบับ

คำสำคัญ: การค้นคืนสารสนเทศ, ประมวลผลภาษาธรรมชาติ, ตัวแบบค้นหา, Doc2Vec

Abstract

This research aims to 1) improve the efficiency of meaningful search subjects from doc2Vec techniques and 2) develop a system for searching for songs from vocals with meaningful search models from Doc2Vec techniques. The data used in the research were Thai lyrics. Research methods consist of data collection, management of Thai lyrics data by natural language processing

methods, and Doc2Vec query model creation. Adjusting the appropriate parameters to enhance the performance of the model Finding the magnitude of the inference vector that is appropriate and results to the search accuracy of the model. Using an optimized model to develop a prototype system for vocal song search.

The findings were as follows: 1) A study to improve the model efficiency found that the size of the questionnaires affecting the search efficiency of the model was 30 percent of the length of the original lyrics to be searched. The size of the vector suitable for modeling in the Thai Folk lyrics dataset is 200. It was found that retaining the stop word in the lyrics resulted in higher efficiency in information retrieval information than eliminating the stop word. 2) From the development of the Doc2Vec song meaning search system for vocal songs. It was found that searchers were able to find songs by inserting vocals through the microphone. The vocals are converted to text with the Web Speech API through the Thai word segmentation process. and create a questionnaire to query with an optimized format. Through testing, we know the appropriate length of the lyrics for the search to get accurate results.

Keywords: Information retrieval, Natural language processing, Search model, Doc2Vec

บทนำ

ในศตวรรษที่ 21 นี้ระบบการสั่งงานด้วยเสียงพูดในคอมพิวเตอร์ อุปกรณ์สมาร์ทโฟน หรืออุปกรณ์แทบเล็ต ได้ถูกพัฒนาการทำงานให้มีประสิทธิภาพสูงขึ้น ผู้ใช้งานสามารถทำกิจกรรมต่าง ๆ ด้วยสมาร์ทโฟนผ่านเสียงพูดได้อย่างรวดเร็ว ความก้าวหน้าของเทคโนโลยีการสั่งงานด้วยเสียงนี้ ส่งผลต่อพฤติกรรมในการสืบค้นข้อมูลของผู้ใช้สมาร์ทโฟนด้วยเช่นกัน โดยจากเดิมที่ต้องพิมพ์คำสำคัญ กลายเป็นการพูดคำ วลี หรือประโยคที่ต้องการทำการสืบค้นเข้ามาแทน แต่อย่างไรก็ตามอุปสรรคสำคัญในการสืบค้นข้อมูลด้วยเสียงของผู้ใช้งาน คือการที่ผู้ใช้ต้องทำการพูดประโยคต่อเนื่องขนาดยาวให้สอดคล้องกับสิ่งที่ต้องการสืบค้น บางครั้งอาจพูดผิดบางคำในประโยคทั้งหมด หรือการประมวลผลที่ผิดพลาดจากโปรแกรมรับคำสั่งเสียงที่เราพบเจอได้อยู่เสมอ ในการศึกษาวิจัยทางด้านการค้นคืนสารสนเทศ (Information Retrieval) ปัจจุบันนิยมนำแบบจำลองปริภูมิเวกเตอร์ (Vector Space Model) มาใช้สำหรับการสร้างตัวแบบสืบค้น ซึ่งจะทำการแปลงเอกสารให้อยู่ในรูปแบบของเมทริกซ์ และคำสอบถาม (Queries) ที่ใช้สืบค้นจะถูกเขียนให้อยู่ในรูปของเวกเตอร์ที่สอดคล้องกับเมทริกซ์เอกสาร [1] ข้อดีของวิธีนี้ ทำให้สามารถสืบค้นเอกสารที่มีความคล้ายคลึงกับคำสอบถามได้ โดยที่คำสอบถามนั้นไม่จำเป็นต้องเกี่ยวข้องกับเอกสารนั้นทั้งหมด แต่อย่างไรก็ตาม การจัดการคำในเอกสารเพื่อนำมาสร้างเมทริกซ์ด้วยเทคนิคถุงคำ (Bag-of-Words) ซึ่งให้ความสนใจเพียงความถี่ของคำที่ปรากฏในเอกสาร แต่ไม่ได้ให้ความสำคัญกับลำดับ และบริบทแวดล้อมของคำภายในเอกสารนั้น ๆ ซึ่งจะส่งผลให้สูญเสียความหมายภายในเอกสารไป และทำให้ยากต่อการใช้สำหรับการค้นหา หรือจำแนกเอกสารในกลุ่มที่คล้ายคลึง หรือมีความใกล้เคียงกันมาก ๆ

Paragraph Vector หรือ Doc2Vec [5] เป็นเทคนิคสำหรับสร้างเวกเตอร์คุณลักษณะ (Feature Vector) ด้วยการแปลงคำ (Word) ให้เป็นตัวเลข (Numerical) ที่อยู่ในรูปของเวกเตอร์ ซึ่งขยายการเรียนรู้การแทนค่าคำศัพท์ด้วยเวกเตอร์ จากระดับของคำไปสู่ระดับประโยค และระดับเอกสาร ด้วยการคงลำดับของคำ ตลอดจนเรียนรู้บริบทแวดล้อมของคำภายในเอกสาร ทำให้เวกเตอร์ที่เกิดขึ้นยังคงความหมายของเอกสารเดิมไว้ ปัจจุบันนี้ Doc2Vec เป็นเทคนิคที่ถูกนำมาใช้ในการสร้างตัวแบบการค้นหาเชิงความหมาย และสามารถจำแนกเอกสารด้วยการวัดความคล้ายเชิงมุม (Cosine Similarity) ได้อย่างมีประสิทธิภาพ [6] [7] [8]

งานวิจัยนี้นำเสนอการเพิ่มประสิทธิภาพของตัวแบบที่ถูกสร้างด้วยเทคนิค Doc2Vec ที่มีความเหมาะสมกับชุดข้อมูล โดยผู้วิจัยนำเทคนิค Doc2Vec มาประมวลผลร่วมกับชุดข้อมูลเนื้อเพลงลูกทุ่งภาษาไทย จำนวน 1,500 บทเพลง สาเหตุเนื่องจากเป็นชุดข้อมูลที่เมื่อแปลงเป็นเวกเตอร์แล้วจะมีความคล้ายคลึง และใกล้เคียงกันเป็นอย่างมาก ซึ่งทำให้ยากต่อการจำแนกความแตกต่างของเอกสารภายในตัวแบบได้ จากนั้นผู้วิจัยได้นำเสนอแนวทางพัฒนาประสิทธิภาพของตัวแบบ ภายใต้ปัญหาการวิจัย คือ 1) ชุดข้อมูลที่นำมาสอนตัวแบบเมื่อกำจัดการหยุดออกไปทำให้บริบทของคำและความหมายเปลี่ยนแปลงไปจะส่งผลต่อประสิทธิภาพของตัวแบบหรือไม่ 2) การสร้างตัวแบบด้วยขนาดของเวกเตอร์ หรือมิติของข้อมูลที่แตกต่างกัน จะส่งผลต่อประสิทธิภาพของตัวแบบหรือไม่ 3) คำสอบถามควรมีขนาดเท่าใดที่จะทำให้ตัวแบบสามารถทำงานด้วยประสิทธิภาพที่เหมาะสม และทำการทดลองเพื่อตอบปัญหาการวิจัยดังกล่าวเพื่อให้ได้ตัวแบบที่มีประสิทธิภาพที่เหมาะสมในการนำไปใช้งานในการประยุกต์สร้างโปรแกรมสำหรับสืบค้นเพลงด้วยเสียงร้องต่อไป

1. วัตถุประสงค์การวิจัย

- 1.1 เพื่อพัฒนาประสิทธิภาพตัวแบบการค้นหาเชิงความหมายจากเทคนิค Doc2Vec
- 1.2 เพื่อพัฒนาระบบการค้นหาเพลงจากเสียงร้องด้วยตัวแบบการค้นหาเชิงความหมายจากเทคนิค

Doc2Vec

2. เอกสารและงานวิจัยที่เกี่ยวข้อง

2.1 Word Embedding

Word Embedding [3] คือการสร้างเวกเตอร์คุณลักษณะ จากคำซึ่งเป็นการแปลงชุดของคำให้เป็นตัวเลขที่ไม่ซ้ำกันโดยการคำนวณความน่าจะเป็นของคำ และบริบทภายในประโยค ซึ่งนำเสนอในรูปของเวกเตอร์ จุดเด่นของเวกเตอร์คุณลักษณะที่ได้จาก Word Embedding นั้น คือสามารถนำไปคำนวณความคล้ายคลึงกับคำอื่น ๆ ในบริบทของคำที่แตกต่างกันได้ ปัจจุบันได้มีการพัฒนาตัวแบบจำลองสำหรับการทำ Word Embedding ออกมาให้ใช้งานอย่างสะดวก เช่น Word2Vec และ Glove โดยงานวิจัยนี้ผู้วิจัยให้ความสนใจแบบจำลอง Word2Vec

Word2Vec [4] ถูกเผยแพร่ในปี 2013 โดย ซึ่งใช้วิธีการโครงข่ายประสาทเทียมเพื่อเรียนรู้การเชื่อมโยงคำจากคลังข้อความขนาดใหญ่ เมื่อได้รับการฝึกอบรมแล้ว โมเดลดังกล่าวสามารถตรวจจับคำที่มีความหมายเหมือนกัน แนะนำคำเพิ่มเติมสำหรับประโยคบางส่วน หรือระดับของความคล้ายคลึงกัน ทางความหมายระหว่างคำที่ถูกแทนด้วยเวกเตอร์ด้วยเทคนิคการวัดค่าความคล้ายเชิงมุม (Cosine Similarity) และ Word2Vec ส่วนมากจะถูกนำไปใช้กับงานทางด้าน Language Modeling ซึ่งเป็นแบบจำลองที่ใช้ในการทำนายคำถัดไปของประโยคว่าควรจะเป็น

เป็นคำใดโดยอาศัยหลักการของ Continuous Bag of Words (CBOW) เพื่อการใช้คำหลาย ๆ คำที่อยู่ต่อเนื่องกันสำหรับการทำนายคำที่อยู่ถัดไป และใช้หลักการ Continuous Skip-Gram สำหรับการใช้คำหนึ่งคำในการทำนายคำอื่น ๆ ที่มีโอกาสเป็นคำถัดไปต่อจากคำปัจจุบัน [6]

2.2 Doc2Vec

Paragraph Vector หรือ Doc2Vec ถูกนำเสนอโดย [5] ซึ่งเป็นส่วนขยายอย่างง่ายของ Word2Vec เพื่อขยายการเรียนรู้ของ Word Embeddings จากระดับของคำไปสู่ระดับประโยค และระดับเอกสาร ซึ่งเป็นการเปลี่ยนจากการเข้ารหัสคำ เป็นการเข้ารหัสข้อความหรือเอกสารทั้งหมด ให้อยู่ในรูปเวกเตอร์คุณลักษณะ โดย Doc2Vec ทำงานด้วยหลักการเดียวกับ Word2Vec ในการเข้ารหัสข้อมูลเพื่อสร้างเวกเตอร์คุณลักษณะ แต่ได้ทำการเพิ่ม Paragraph ID เข้ามาในกระบวนการสร้าง โดย Paragraph ID ถูกสร้างมาจากตัวเอกสาร และทำให้มีเลขเฉพาะตัวในการนำไปใช้วิเคราะห์ จุดเด่นของ Doc2Vec คือความสามารถในการแก้ปัญหาในการจำกัดจำนวนคุณลักษณะที่ใช้สำหรับการเรียนรู้ของแบบจำลอง ดังเช่น Bag-of-Word และ TF-IDF ที่จะต้องกำหนดจำนวนคำศัพท์ที่จะนำมาสร้างเวกเตอร์คุณลักษณะ โดย Doc2Vec นี้ประสบความสำเร็จอย่างมากในการนำไปประยุกต์ใช้งานทางด้าน การหาความคล้ายคลึงของเอกสาร การจำแนกประเภทของเอกสาร และการประยุกต์เพื่อใช้งานด้านเครื่องจักรแปลภาษา

2.3 งานวิจัยที่เกี่ยวข้อง

Le และ Mikolov [5] นำเสนอเทคนิค Paragraph Vector หรือ Doc2Vec ซึ่งเป็นอัลกอริทึมการเรียนรู้แบบไม่มีผู้สอนซึ่งใช้เพื่อการเรียนรู้การแทนค่าเวกเตอร์สำหรับข้อความที่มีความยาวผันแปรได้ เช่น ประโยคและเอกสาร การแทนเอกสารด้วยเวกเตอร์นั้น ทำให้เกิดการเรียนรู้ที่สามารถใช้สำหรับการทำนายคำจากบริบทแวดล้อมได้ การบ่งชี้ถึงประสิทธิภาพของ Doc2Vec ผู้วิจัยทั้งสองได้ทำการทดลองเกี่ยวกับการจำแนกข้อความจากการทดลองแสดงให้เห็นว่า Doc2Vec มีประสิทธิภาพในการจับความหมายของเอกสารด้วยการเรียนรู้ของเครื่อง ซึ่งเป็นการลบข้อเสียของโมเดลการทำงานแบบถูคำ (Bag-of-Word)

Andrew และคณะ [6] ได้ทำการทดสอบเพื่อวัดประสิทธิภาพของ Doc2Vec ด้วยการเปรียบเทียบกับวิธีการจำแนกเอกสารแบบอื่น คือ Latent Dirichlet Allocation (LDA) และ Bag-of-Word (BOW) รู้โดยใช้ชุดข้อมูลจาก Wikipedia และ arXiv สำหรับการสอนตัวแบบให้เรียนรู้ ผลการทดสอบบ่งชี้ว่า Doc2Vec มีประสิทธิภาพสูงกว่าวิธีการที่นำมาเปรียบเทียบในด้านของการวัดค่าความคล้ายคลึงเชิงความหมายของเอกสาร และการทดสอบนี้ยังบ่งชี้ว่า ขนาดเวกเตอร์ หรือขนาดมิติข้อมูลที่ใช้สร้างตัวแบบนั้นส่งผลต่อประสิทธิภาพการทำงานของตัวแบบ โดยตัวแบบจะทำงานได้อย่างมีประสิทธิภาพสูงสุดที่ขนาดของเวกเตอร์ค่าใดค่าหนึ่งที่มีความเหมาะสม นอกจากนี้การทดลองของผู้วิจัยแสดงให้เห็นว่าเมื่อนำข้อมูลจากเวกเตอร์มาแนะนำเสนอด้วย T-SNE ข้อมูลที่อยู่ในหมวดหมู่เดียวกันจะมีความสัมพันธ์กันและเกาะกลุ่มอยู่ด้วยกันอย่างใกล้ชิด นอกจากนี้ Lau และ Baldwin [7] พบว่าการสอนตัวแบบ Doc2Vec จากไลบรารี Gensim ด้วย Pre-Trained Word Embedding ที่ถูกสอนด้วยคลังข้อมูล WIKI และ AP-NEWS สามารถเพิ่มประสิทธิภาพของ ตัวแบบที่ใช้ดำเนินงานปัจจุบันได้ตลอดจนได้แนะนำการกำหนดค่าไฮเปอร์พารามิเตอร์ที่เหมาะสมและส่งผลต่อประสิทธิภาพของ Doc2Vec

สำหรับงานวิจัยด้านการประยุกต์ใช้งาน Word Embedding ประกอบด้วยงานวิจัยที่นำ Word2Vec มาใช้ในการแปลงชุดคำให้เป็นเวกเตอร์ที่สามารถบ่งชี้ความคล้ายเชิงความหมายระหว่างกันเพื่อเพิ่มความแม่นยำ ของ

ตัวแบบจำแนกโดย บุขบงก์ คชินทรโรจน์ และคณะ [2] ใช้ Word2Vec เป็นหนึ่งในวิธีการแปลงคำให้เป็นเวกเตอร์ ก่อนนำไปสร้างตัวจำแนกด้วยวิธี LDA, Kaothanthong และคณะ [8] สร้างตัวแบบเพื่อจำแนกหัวข้อข่าวหลอกลวง ภาษาไทย โดยสร้างเวกเตอร์คุณลักษณะจาก 2 วิธี คือ 1) เวกเตอร์คุณลักษณะที่เรียกว่า Headline2Vec ซึ่งได้ Concatenation Layer ซึ่งเป็นชั้นสุดท้ายของโครงข่ายประสาทเทียมที่ใช้ใน Word2Vec และ 2) เวกเตอร์คุณลักษณะที่เกิดจากการคำนวณด้วยวิธี TF-IDF จากนั้นนำเวกเตอร์คุณลักษณะจากวิธีทั้งสองมาสร้างตัวแบบจำแนกด้วยอัลกอริทึม Support Vector Machine, naïve Bayes และ Multilayer Perceptron โดยผลการเปรียบเทียบพบว่าเวกเตอร์คุณลักษณะ Headline2Vec สามารถนำไปสร้างตัวแบบที่มีประสิทธิภาพในการจำแนกได้ดีกว่า TF-IDF ซึ่งนำไปสู่ข้อสรุปของการทดลองว่า ความสัมพันธ์ระหว่างคำในเวกเตอร์คุณลักษณะมีผลต่อประสิทธิภาพการทำงานของตัวแบบจำแนก นอกจากนี้ในการทดลองผู้วิจัยได้นำเสนอถึงการปรับค่าไฮเปอร์พารามิเตอร์ที่เหมาะสมในการสร้างเวกเตอร์คุณลักษณะ ซึ่งบ่งชี้ว่าจำนวนมิติข้อมูลหรือขนาดของเวกเตอร์ค่าใดค่าหนึ่งที่มีความเหมาะสมจะส่งผลต่อประสิทธิภาพสูงสุดของตัวแบบ

สำหรับงานวิจัยที่นำ Doc2Vec มาประยุกต์ใช้งานประกอบด้วย Alshammeria และคณะ [9] ได้นำ Doc2Vec มาใช้สำหรับวิเคราะห์ความคล้ายเชิงความหมายของเอกสารคัมภีร์กุรอ่านต้นฉบับ และฉบับแปลภาษาอังกฤษ Budiarto และคณะ [10] ใช้ Doc2Vec ร่วมกับอัลกอริทึมการจัดกลุ่มข้อมูลเพื่อสร้างตัวแบบสำหรับจัดกลุ่มหัวข้อข่าวภาษาอินโดนีเซีย Patra และคณะ [11] พัฒนาระบบแนะนำเพลงภาษาฮินดี (Hindi) จากความคล้ายของเนื้อเพลง โดยใช้บทเพลงฮินดีจำนวน 31,171 เพลง ด้วยการรวบรวมจากเว็บไซต์ โดยได้ทำการทดสอบเปรียบเทียบประสิทธิภาพตัวแบบที่สร้างด้วย Doc2Vec และ Self-Organizing Feature Maps (SOFMs) โดยในข้อสรุปการวิจัยผู้วิจัยได้สรุปผลว่าตัวแบบที่สร้างด้วย Doc2Vec มีประสิทธิภาพ ไม่สูงนักอาจเนื่องมาจาก Doc2Vec ต้องการคลังข้อมูลขนาดใหญ่ในการสอนตัวแบบให้เรียนรู้ จำนวนเพลงภาษาฮินดีที่นำมาใช้สอนอาจมีปริมาณไม่เพียงพอ และ Guli และ Tiwari [12] ได้พยายามนำเอาเทคนิค Doc2Vec เพื่อใช้สำหรับการพัฒนาระบบแนะนำเพลงด้วยเสียงพูดด้วยเช่นกันโดยคาดหวังว่าจะช่วยให้ผู้ใช้งานสามารถเข้าถึงเพลงที่ต้องการได้สะดวกยิ่งขึ้น

วิธีดำเนินการวิจัย

1. เครื่องมือการวิจัย

ในกระบวนการพัฒนาประสิทธิภาพตัวแบบจากเทคนิค Doc2Vec เพื่อการประยุกต์ใช้สำหรับการสืบค้นเพลงด้วยเสียงร้อง ผู้วิจัยได้ใช้เครื่องมือต่าง ๆ ดังนี้คือ 1) ภาษาไพธอน เวอร์ชัน 3.7.6 ในการพัฒนาโปรแกรม 2) ไลบรารี PythaiNLP เวอร์ชัน 2.2.3 สำหรับใช้ตัดคำภาษาไทย 3) ไลบรารี Pandas เวอร์ชัน 1.0.1 สำหรับจัดการและประมวลผลข้อมูล 4) คลาส Doc2Vec จากไลบรารี Gensim เวอร์ชัน 3.8.3 สำหรับใช้แปลงเอกสารให้เป็นเวกเตอร์คุณลักษณะ 5) คลาส TSNE จากไลบรารี scikit-learn เวอร์ชัน 0.22.1 ซึ่งใช้สำหรับลดมิติข้อมูลแบบไม่เชิงเส้น 6) Web Speech API สำหรับการแปลงเสียงพูด หรือการร้องเพลงให้กลายเป็นข้อความภาษาไทยเพื่อการสืบค้นข้อมูลเพลง 7) Django Web Framework เวอร์ชัน 3.6.8 สำหรับการพัฒนาโปรแกรมประยุกต์เพื่อการสืบค้นเพลงจากเสียงร้องในลักษณะของเว็บแอปพลิเคชัน

2. กลุ่มเป้าหมาย

กลุ่มเป้าหมายในการทดสอบโปรแกรม และรับข้อมูลป้อนกลับจากการทดลองใช้งานโปรแกรมเพื่อนำกลับไปพัฒนาโปรแกรมให้ดียิ่งขึ้น เป็นนักศึกษาชั้นปีที่ 4 สาขาวิทยาการคอมพิวเตอร์ มหาวิทยาลัยราชภัฏ ศรีสะเกษ จำนวน 20 คน ซึ่งถูกเลือกแบบเจาะจง โดยเป็นนักศึกษาที่ลงทะเบียนเรียนในรายวิชาการสืบค้นข้อมูลสารสนเทศ ปีการศึกษา 1/2563

3. ขั้นตอนการดำเนินการวิจัย

3.1 การจัดเก็บและรวบรวมข้อมูล

ในการนำเนื้อเพลงมาสร้างเป็นตัวแบบ คำศัพท์ต่าง ๆ ที่อยู่ในเนื้อเพลงล้วนมีความหมายต่อการพิจารณาเพื่อสร้างเป็นตัวแบบ เพื่อไม่ให้เกิดความแตกต่างของคำศัพท์ที่อยู่ในเนื้อเพลงที่จะนำมาสร้างเป็นตัวแบบมากเกินไป ดังนั้นการทดลองในงานวิจัยนี้ ผู้วิจัยจึงได้พิจารณาเฉพาะบทเพลง “ลูกทุ่ง” มาให้สำหรับการสร้างตัวแบบ เพราะคำศัพท์และวิธีการเรียบเรียงประโยคมีความใกล้เคียงและคล้ายคลึงกันในแต่ละบทเพลง และผู้วิจัยได้ดำเนินการรวบรวมข้อมูลเนื้อเพลงภาษาไทย จากเว็บไซต์ youtube.com, sanook.com, kapook.com, joox.com โดยเนื้อเพลงลูกทุ่งที่ผู้วิจัยได้รวบรวมและใช้สำหรับการศึกษาวิจัยครั้งนี้มีจำนวนทั้งสิ้น 1,500 เพลง ซึ่งเป็นเพลงที่ออกเผยแพร่อยู่ในช่วงปี ค.ศ. 2000 - 2021

3.2 การเตรียมข้อมูลและสร้างตัวแบบสืบค้นเชิงความหมาย

การเตรียมข้อมูลเนื้อเพลงภาษาไทยในงานวิจัยนี้ ผู้วิจัยได้ใช้กระบวนการของการประมวลผลภาษาธรรมชาติสำหรับจัดเตรียมข้อมูล โดยสามารถนำเสนอการดำเนินงานออกเป็นขั้นตอนดังต่อไปนี้

ขั้นที่ 1 กระบวนการตัดคำภาษาไทย (Thai Words Segmentation) โดยใช้ไลบรารี PyThaiNLP ด้วยเทคนิคการตัดคำภาษาไทยแบบ Maximum Matching Algorithm

ขั้นที่ 2 กระบวนการกำจัดคำหยุด (Stop Words Removal) ซึ่งเป็นการกำจัดคำที่ไม่สื่อความหมายที่ปรากฏอยู่ในเนื้อเพลงออกไป โดยผู้วิจัยใช้รายการคำหยุดจากไลบรารี PyThaiNLP จำนวน 1030 คำ เป็นตัวเปรียบเทียบเพื่อคัดกรอง โดยผู้วิจัยได้ทำการจัดเตรียมข้อมูลออกเป็นทั้งหมด 2 ชุด คือ ชุดข้อมูลเนื้อเพลงที่กำจัดคำหยุดออกไป และชุดข้อมูลเนื้อเพลงที่ยังไม่มีการกำจัดคำหยุด

ขั้นที่ 3 สร้างตัวแบบสืบค้นเชิงความหมายด้วยเทคนิค Doc2Vec จากไลบรารี Gensim โดยกำหนดค่าไฮเปอร์พารามิเตอร์ [7] ดังนี้ min_count=2 และ epochs=40 และสำหรับ vector_size เป็นการกำหนดมิติเวกเตอร์คุณลักษณะ (Dimensionality of The Feature Vectors) โดยผู้วิจัยต้องการทราบว่าความแตกต่างนี้จะส่งผลต่อความแม่นยำในการสืบค้นของตัวแบบหรือไม่ ผู้วิจัยได้ทำการสร้างตัวแบบที่หลากหลายโดยกำหนดขนาดของ vector_size ที่แตกต่างกันตั้งแต่ 500, 400, 300, 200, 100 และ 50

3.3 การหาขนาดที่เหมาะสมของข้อมูลป้อนเข้าเพื่อการค้นหา

การค้นคืนสารสนเทศเพื่อค้นหาเอกสารที่มีคล้ายคลึงกันนั้น ประเด็นในเรื่องของขนาดข้อมูลป้อนเข้า (Input) ยังเป็นอีกปัญหาที่น่าสนใจในการพัฒนาประสิทธิภาพความแม่นยำของตัวแบบการสืบค้นอยู่เสมอ เช่น เมื่อผู้สืบค้นต้องค้นหาเพื่อให้ได้เอกสารที่ต้องการนั้น ผู้สืบค้นจะต้องป้อนข้อมูลขนาดความยาวเท่าใด เพื่อให้ผลของการค้นหานั้นออกมาแม่นยำมากที่สุด และสำหรับปัญหาของการวิจัยนี้เป็นการค้นหาคำเพลงด้วยเสียงร้องนั้น ผู้วิจัยต้องการทราบคำตอบว่า ผู้สืบค้นจะต้องทำการร้องเนื้อเพลงเป็นความยาวเท่าใดเพื่อให้การสืบค้นบทเพลงมีความแม่นยำมากที่สุด

ในการทดลองนี้ผู้วิจัยได้ออกแบบกระบวนการทดสอบความแม่นยำในการสืบค้นของตัวแบบด้วยการสุ่มเลือกบทเพลงขึ้นมาครั้งละ 100 บทเพลง จากรายการบทเพลงทั้งหมดที่ใช้ในการสอนตัวแบบ และทำการกำหนดค่าสัดส่วนความยาวของคำสอบถามสำหรับสร้างเป็นเวกเตอร์อนุมาน (Inferred Vector) เพื่อการสืบค้น โดยกำหนดความยาวเป็นร้อยละ 50, 40, 30, 20, 19, 18, 17, 16, 15, 10, 7, 5 ของเนื้อเพลงต้นฉบับ ในการปรับค่าความยาวของเนื้อเพลงป้อนเข้าให้เหลือสัดส่วนของความยาวตามที่กำหนด ผู้วิจัยได้ทำการสุ่มเลือกตัดเนื้อเพลงบางส่วนออกไป แต่ยังคงลำดับก่อน-หลังของคำภายในเนื้อเพลงไว้ เพื่อให้คงความหมายเดิมของเพลงให้มากที่สุด และนำเอาท่อนเพลงที่ประมวลผลแล้วไปทำการทดสอบความแม่นยำในการสืบค้นกับทุก ๆ ตัวแบบที่ถูกสร้างขึ้น ด้วยมิติของข้อมูลที่แตกต่างกันต่อไป

3.4 การวัดประสิทธิภาพตัวแบบ

ในการวัดประสิทธิภาพการสืบค้นของของตัวแบบค้นหาเชิงความหมายจากเทคนิค Doc2Vec นี้ผู้วิจัยใช้ค่าความแม่นยำ (Accuracy) สำหรับการบ่งชี้ประสิทธิภาพของตัวแบบ โดยวิธีการคำนวณค่าความแม่นยำที่ใช้แสดงได้ดังต่อไปนี้

$$\text{Accuracy \%} = 100 - \text{Error \%}$$

โดยที่

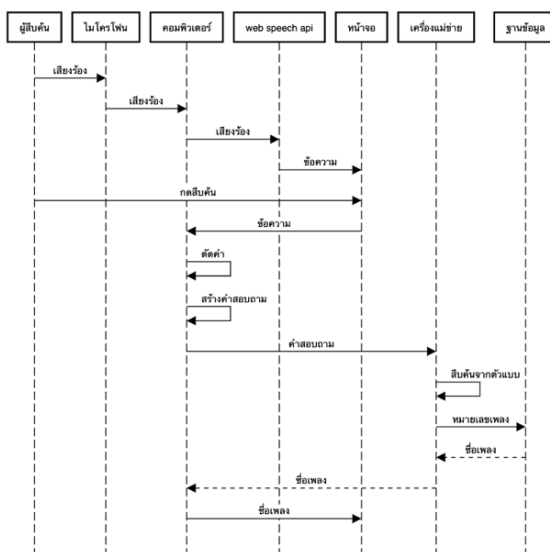
$$\text{Relative error} = \left| \frac{X_{mea} - X_t}{X_t} \right|$$

เมื่อ X_{mea} คือ ค่าที่ได้จากการวัดจริง (Measure Value)

X_t คือ ค่าจริง (True Value)

3.5 การพัฒนาโปรแกรมประยุกต์เพื่อการสืบค้นเพลงจากเสียงร้อง

เมื่อผู้วิจัยทำการทดสอบเพื่อหาตัวแบบสืบค้นเชิงความหมายที่มีประสิทธิภาพเหมาะสมที่สุดแล้ว จะทำการนำแบบจำลองนั้นไปใช้สำหรับการพัฒนาโปรแกรมประยุกต์ต่อไป สำหรับการรับเสียงพูดหรือการร้องจากผู้ใช้งาน ผู้วิจัยใช้ Web Speech API เพื่อการประมวลผลเสียงพูดให้เป็นตัวอักษรภาษาไทย โดยการออกแบบและพัฒนาโปรแกรมประยุกต์ดังกล่าวนี้ สามารถแสดงได้ดังภาพที่ 1



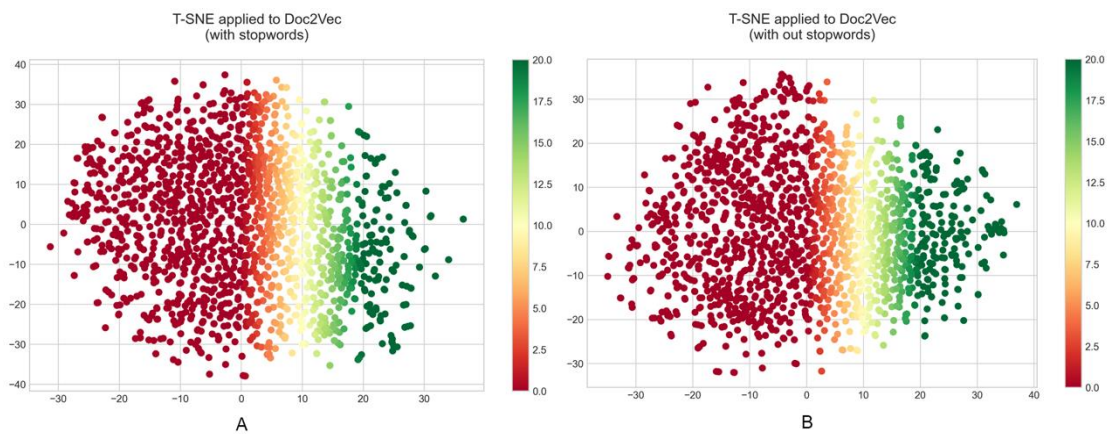
ภาพที่ 1 แสดงลำดับการทำงานของโปรแกรมประยุกต์เพื่อการสืบค้นเพลงจากเสียงร้อง

ผลการวิจัย

1. ผลการพัฒนาประสิทธิภาพตัวแบบค้นหาเชิงความหมายจากเทคนิค Doc2Vec

ผู้วิจัยได้ดำเนินการพัฒนาประสิทธิภาพของตัวแบบภายใต้ปัญหาของการวิจัยดังนี้ 1) ชุดข้อมูลที่นำมาสอนตัวแบบเมื่อกำจัดคำหยาบออกไปทำให้บริบทของคำและความหมายเปลี่ยนแปลงไปจะส่งผลต่อประสิทธิภาพของตัวแบบหรือไม่ 2) การสร้างตัวแบบด้วยขนาดของเวกเตอร์ หรือมิติของข้อมูลที่แตกต่างกัน จะส่งผลต่อประสิทธิภาพของตัวแบบหรือไม่ 3) คำสอบถามควรมีขนาดเท่าใดที่จะทำให้ตัวแบบสามารถทำงานด้วยประสิทธิภาพที่เหมาะสม การดำเนินงานเพื่อหาคำตอบของปัญหาการวิจัยนั้น ผู้วิจัยได้จัดเตรียมข้อมูลเนื้อเพลงจำนวน 1,500 บทเพลงที่ผ่านกระบวนการตัดคำภาษาไทย โดยแบ่งออกเป็น 2 ชุด ประกอบด้วย ชุดข้อมูลต้นฉบับ และชุดข้อมูลที่กำจัดคำหยาบออก จากนั้นนำข้อมูลทั้งสองชุดไปสร้างตัวแบบด้วยเทคนิค Doc2Vec และสร้างความหลากหลายของตัวแบบด้วยการกำหนดขนาดของเวกเตอร์ที่ต่างกัน

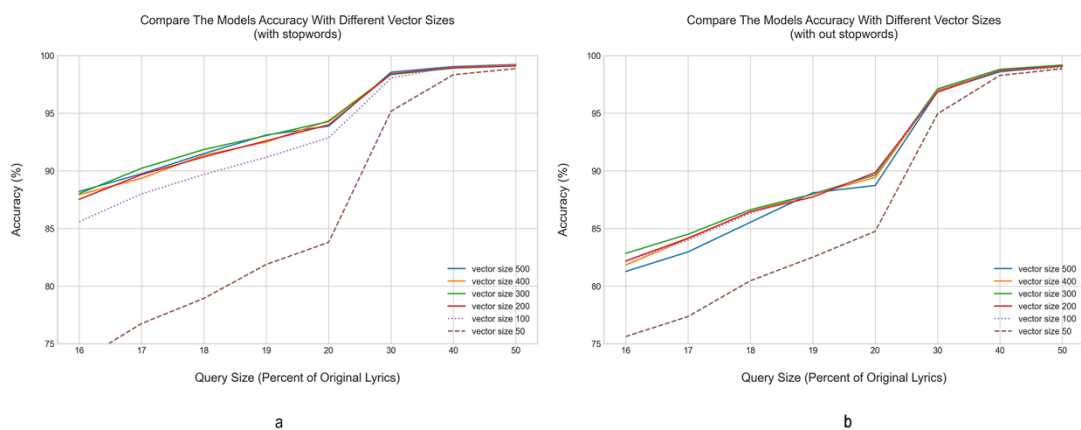
โดยเวกเตอร์ซึ่งเป็นตัวแทนของบทเพลงแต่ละเพลงซึ่งได้จากเทคนิค Doc2Vec เป็นชุดข้อมูลที่มีจำนวนมิติมาก (High Dimensional Data) และเพื่อแสดงให้เห็นถึงความสัมพันธ์ของเวกเตอร์เนื้อเพลงทั้ง 1,500 เพลงในตัวแบบที่ถูกขึ้น ผู้วิจัยใช้ T-SNE (t-distribution Stochastic Neighbor Embedding) ซึ่งเป็นเทคนิคที่ใช้ลดมิติของข้อมูลที่มีจำนวนมิติมาก ให้สามารถแปลงเวกเตอร์เอกสารที่มีขนาดมิติข้อมูลเท่ากับ 200 ให้สามารถนำเสนอได้ในระนาบ 2 มิติ โดยใช้แผนภาพการกระจาย (Scatter Plot) ดังภาพที่ 2



ภาพที่ 2 แผนภาพการกระจายแสดงความสัมพันธ์ของเนื้อเพลงภายในตัวแบบ

จากภาพที่ 2 แสดงให้เห็นว่าลักษณะความสัมพันธ์ของเนื้อเพลงจำนวน 1,500 เพลงที่นำมาใช้สร้างโมเดล เนื้อเพลงมีความคล้ายคลึงกัน โดยข้อมูลมีการเกาะกลุ่มกันอย่างใกล้ชิดเนื่องจากเนื้อเพลงที่นำมาใช้ทดสอบในการวิจัยนี้อยู่ในกลุ่มเพลง “ลูกทุ่ง” เหมือนกัน เมื่อเปรียบเทียบการกระจายตัวของข้อมูล พบว่าชุดข้อมูลที่ยังคงคำหยุดไว้ในเนื้อเพลงจะมีความสัมพันธ์ใกล้ชิดกันมาก และมีการกระจายตัวของข้อมูลน้อยกว่าการนำเอาคำหยุดออกจากเนื้อเพลง

จากนั้นผู้วิจัยได้ดำเนินการทดลองเพื่อตอบปัญหาการวิจัยที่ตั้งไว้ทั้ง 3 ประเด็น ซึ่งจะทำได้ทำให้สามารถค้นพบตัวแบบที่มีประสิทธิภาพ และเหมาะสมต่อการนำไปใช้ในการพัฒนาโปรแกรมประยุกต์ที่ต้องการต่อไปได้ ซึ่งสามารถแสดงผลของการทดลองได้ดังภาพที่ 3



ภาพที่ 3 ผลการทดลองเพื่อเปรียบเทียบประสิทธิภาพของตัวแบบ

โดยจากภาพที่ 3 (a) แสดงความสัมพันธ์ระหว่างขนาดของคำสอบถาม และร้อยละของความแม่นยำในการสืบค้น กรณีที่ยังคงคำหยุดไว้ในเนื้อเพลง เมื่อพิจารณาขนาดของเวกเตอร์ที่ใช้สร้างตัวแบบตั้งแต่ 50 ถึง 500 จะเห็นได้ว่าเมื่อเพิ่มขนาดร้อยละของคำสอบถามจะพบว่าความแม่นยำในการค้นหาคำของตัวแบบมีค่าสูงขึ้น โดยพบว่าการใช้ขนาดของคำสอบถามที่ร้อยละ 30 ส่งผลให้ความแม่นยำในการค้นหามีค่ามากกว่าร้อยละ 95 ในทุก ๆ ตัว

แบบที่มีขนาดเวกเตอร์ต่างกัน และหากพิจารณาถึงการเพิ่มขนาดร้อยละของคำสอบถามเป็น 40 – 50 พบว่าส่งผลให้ผลประสิทธิภาพการค้นหาเพิ่มขึ้นเพียงเล็กน้อย และเมื่อพิจารณาถึงขนาดของเวกเตอร์ที่นำมาสร้างตัวแบบที่มีความเหมาะสมกับชุดข้อมูลเนื้อเพลงลูกทุ่งนี้ จะพบว่าช่วงของขนาดเวกเตอร์ที่มีความเหมาะสมในการใช้สร้างตัวแบบอยู่ระหว่าง 200 – 500 และเมื่อพิจารณาจากค่าร้อยละของความแม่นยำพบว่า ตัวแบบที่สร้างโดยใช้ขนาดเวกเตอร์เท่ากับ 200 ให้ค่าความแม่นยำในการค้นหาที่ร้อยละ 98.41, 98.86, 98.97 ที่ขนาดคำสอบถามร้อยละ 30, 40 และ 50 ตามลำดับ ดังแสดงในตารางที่ 1 โดยชุดข้อมูล 1 หมายถึงชุด เนื้อเพลงที่คงคำหยุดไว้ และชุดข้อมูล 2 หมายถึงเนื้อเพลงที่กำลังจัดคำหยุดออกไป

ตารางที่ 1 ผลการทดลองวัดประสิทธิภาพความแม่นยำของตัวแบบ

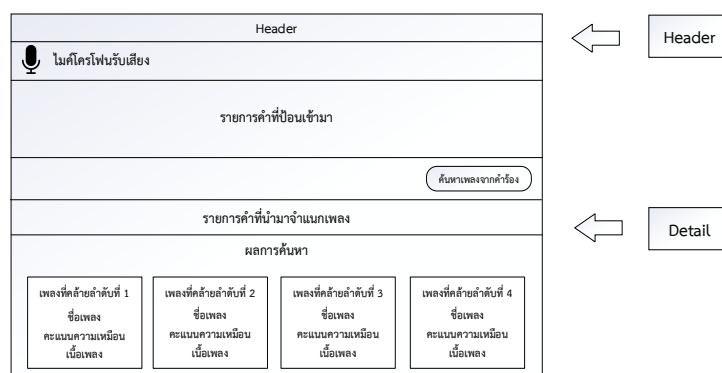
Vector size	ขนาดคำสอบถาม (%)															
	16		17		18		19		20		30		40		50	
	1	2	1	2	1	2	1	2	1	2	1	2	1	2	1	2
500	88.21	81.27	89.75	82.97	91.50	85.54	93.11	88.10	93.88	88.73	98.56	96.89	99.05	98.59	99.23	99.12
400	87.93	81.83	89.36	84.12	91.37	86.47	92.50	87.96	94.35	89.42	98.38	96.97	98.98	98.69	99.17	99.07
300	88.03	82.84	90.22	84.50	91.85	86.63	93.06	87.99	94.28	89.64	98.35	97.11	98.91	98.81	99.10	99.18
200	88.01	82.18	90.12	84.59	91.29	87.46	93.20	88.26	94.00	90.27	98.41	97.18	98.86	98.86	98.97	99.09
100	85.59	82.13	88.00	83.99	89.69	86.32	91.18	88.00	92.87	89.69	98.04	97.02	99.06	98.63	99.20	98.96
50	73.45	75.62	76.75	77.36	78.94	80.47	81.88	82.51	83.80	84.75	95.17	94.95	98.34	98.28	98.87	98.85

จากตารางที่ 3 (b) แสดงความสัมพันธ์ระหว่างขนาดของคำสอบถาม และร้อยละของความแม่นยำในการสืบค้น กรณีที่กำลังจัดคำหยุดออกจากเนื้อเพลง พบว่าแนวโน้มของประสิทธิภาพของตัวแบบเป็นไปในทิศทางเดียวกับกรณีที่ไม่วางจัดคำหยุด การใช้ขนาดคำสอบถามที่ร้อยละ 30 ขึ้นไปส่งผลให้ความแม่นยำในการค้นหามีค่ามากกว่าร้อยละ 95 เช่นกัน แต่มีค่าร้อยละของความแม่นยำน้อยกว่าการคงคำหยุดไว้ในเนื้อเพลงในทุก ๆ ตัวแบบที่มีขนาดเวกเตอร์แตกต่างกัน

2. ผลการพัฒนาระบบการค้นหาเพลงจากเสียงร้องด้วยตัวแบบการค้นหาเชิงความหมายจากเทคนิค Doc2Vec

2.1 ผลการออกแบบองค์ประกอบของโปรแกรมประยุกต์เพื่อสืบค้นเพลงด้วยเสียงร้อง

จากกระบวนการออกแบบระบบด้วยการเก็บข้อมูล และข้อคิดเห็นจากผู้ใช้งานพบว่า โปรแกรมประยุกต์เพื่อการสืบค้นเพลงด้วยเสียงร้องที่มีความเหมาะสมต่อการใช้งานประกอบด้วย องค์ประกอบ 3 ส่วน ได้แก่ 1) ส่วนหัวของเว็บไซต์ 2) ส่วนรับข้อมูลป้อนเข้าจากผู้ใช้งานด้วยเสียงร้อง และ 3) ส่วนสำหรับการแสดงผลการสืบค้น ดังแสดงดังภาพที่ 4



ภาพที่ 4 องค์ประกอบหน้าจอของโปรแกรมสืบค้นเพลงด้วยเสียงร้อง

2.2 ผลการพัฒนาโปรแกรมประยุกต์เพื่อสืบค้นเพลงด้วยเสียงร้อง

ผู้วิจัยทำการพัฒนาโปรแกรมประยุกต์เพื่อสืบค้นเพลงด้วยเสียงร้องในรูปแบบของเว็บแอปพลิเคชันตามที่ได้ออกแบบไว้ ผู้วิจัยใช้ Web Speech API ทำหน้าเป็นไคลเอนต์ (Client) รับข้อมูลเสียงพูดและแปลงเป็นข้อความภาษาไทย และส่งให้ระบบแบ็คเอนด์ (Back-End) ด้วยระบบเว็บเซอร์วิส (Rest API Web Service) ซึ่งผู้วิจัยพัฒนาโดยใช้ Django และ Rest Framework เมื่อฝั่งแบ็คเอนด์ได้รับข้อมูลที่เป็นข้อความภาษาไทย จะนำไปประมวลด้วยการตัดคำภาษาไทย และทำการแปลงให้เป็นคำสอบถามที่อยู่ในรูปของเวกเตอร์และส่งให้ตัวแบบค้นหาเพลงที่มีเนื้อเพลงคล้ายกับคำสอบถามมากที่สุด และส่งคำตอบที่เป็นชื่อเพลงให้แก่มือสืบค้น โดยสามารถแสดงหน้าจอการทำงานของโปรแกรมที่เป็นผลของการพัฒนาได้ดังภาพที่ 5

ค้นเพลงลูกทุ่ง

ขอที่จะร่ำเฝ้านับวันกาลให้เจอคนดีให้สาวไกลบ้านคนดีที่มีคนคอยตอบเหงาจนขุ่นมืองใหญ่ขอใจข้ามถิ่นกันหนาว
น้ำพาในทุ่งหรือจวาวให้สาวตอบทงภูจำนึ่งใจ

🎤

ค้นหาเพลง

ผลการค้นหาเพลง			
ป้ายกำกับ	ชื่อเพลง	ศิลปิน	คะแนนความคล้าย
MOST	ขอใจกันหนาว	ค่าย อรทัย	0.7895627617835999
SECOUND	ผ้าขาวบนผ้าขาว	ไมค์ ภิรมย์พร	0.5311357378959656
THIRD	ขอใจไว้เรียงจันทร์	อัมพร แพรนเพชร	0.46886077523231506
FOUTH	สาวคำบกทุ่ง	ศิริพร อำไพพงษ์	0.4575023651123047
FIFTH	หนึ่งก่าลิ้ง	ฝน ธนสุนทร	0.4537308812141485
MEDIAN	มีทองท่วมหัว ไม่มีม่วงก็ได้	จีระ ชาร์ชาม	0.20947681367397308
LEAST	กาบพาไร่เลย	ศิริพร อำไพพงษ์	-0.07151535898447037

ภาพที่ 5 หน้าจอการทำงานของโปรแกรมสืบค้นเพลงด้วยเสียงร้อง

หลังจากนำโปรแกรมสืบค้นเพลงด้วยเสียงร้องไปให้กลุ่มเป้าหมายจำนวน 20 คนได้ทำการทดลองใช้งาน และประเมินประสิทธิภาพการทำงานของโปรแกรม ทั้งหมด 3 ด้าน ประกอบด้วย ด้านการใช้ประโยชน์ ด้านการใช้งานโปรแกรม ด้านความถูกต้องในการทำงานของโปรแกรมจากผลการประเมินพบว่า มีประสิทธิภาพโดยรวมอยู่ในระดับมาก ($\bar{x}=4.11$, $SD.=0.66$) เมื่อพิจารณาเป็นรายด้านพบว่า ด้านการใช้ประโยชน์ของโปรแกรมมีประสิทธิภาพอยู่ในระดับมาก ($\bar{x}=4.17$, $SD.=0.69$) ด้านการใช้งานโปรแกรมมีประสิทธิภาพอยู่ในระดับมาก ($\bar{x}=4.13$, $SD.=0.65$) และด้านความถูกต้องในการทำงานของโปรแกรมมีประสิทธิภาพอยู่ในระดับมาก ($\bar{x}=3.98$, $SD.=0.65$)

อภิปรายผลการวิจัย

1. การพัฒนาประสิทธิภาพของตัวแบบค้นหาเชิงความหมายจากเทคนิค Doc2Vec สามารถอภิปรายผลการวิจัยได้ดังนี้ จากกระบวนการศึกษาผู้วิจัยได้ใช้ T-SNE นำเสนอความสัมพันธ์ของข้อมูลหลังจากนำไปสร้างเป็นเวกเตอร์คุณลักษณะ พบว่าข้อมูลมีความสัมพันธ์ และมีการเกาะกลุ่มกันอย่างหนาแน่น เนื่องจากเป็นข้อมูลเพลงที่อยู่ในหมวดหมู่เดียวกัน การใช้คำ และการจัดวางบริบทแวดล้อมของคำมีความคล้ายคลึงกันเป็นอย่างมาก ซึ่งสอดคล้องกับ Andrew และคณะ [6] ที่นำเสนอว่าข้อมูลในหมวดหมู่เดียวกันจะมีความใกล้ชิดกัน นอกจากนี้ผู้วิจัยพิจารณาถึงการค้นหาค่าไฮเปอร์พารามิเตอร์ที่เหมาะสม โดยการทดสอบพบว่ามิติข้อมูล หรือขนาดเวกเตอร์ ที่ใช้สร้างตัวแบบด้วย Doc2Vec สำหรับชุดข้อมูลเพลงลูกทุ่งภาษาไทยมีขนาดอยู่ที่ 200 และถึงแม้ว่าจะเพิ่มขนาดของเวกเตอร์ให้มากขึ้น แต่ก็ไม่ส่งผลต่อประสิทธิภาพของตัวแบบให้สูงขึ้นแต่อย่างใด สอดคล้องกับการทดลองของ Andrew และคณะ [6] และ Kaothanthong และคณะ [8] ที่บ่งชี้ว่าประสิทธิภาพของตัวแบบ Doc2Vec จะสูงสุดที่ขนาดของเวกเตอร์ค่าใดค่าหนึ่ง และผลการทดลองบ่งชี้ให้เห็นว่าตัวแบบที่ถูกสอนด้วยชุดข้อมูลที่ไม่จำกัดคำหยุด

มีประสิทธิภาพดีกว่าตัวแบบที่ถูกสอนด้วยชุดข้อมูลที่กำหนดค่าหยุดออกไป ซึ่งแสดงให้เห็นว่าการกำหนดค่าหยุดออกไปนั้นทำให้บริบทของคำ ตลอดจนความหมายของเอกสารเปลี่ยนแปลงไป ซึ่งทำให้เทคนิค Doc2Vec ไม่สามารถเรียนรู้ความสัมพันธ์ระหว่างคำได้ดีเท่าที่ควร ผลการศึกษานี้สอดคล้องกับงานวิจัยของ Kaothanthong และคณะ [8] ว่าความสัมพันธ์ระหว่างคำในเวกเตอร์คุณลักษณะมีผลต่อการทำงานของตัวแบบเป็นอย่างมาก และการทดสอบเพื่อหาขนาดคำสอบถามที่เหมาะสมในการนำไปสร้างเวกเตอร์อนุมาณเพื่อส่งให้ตัวแบบค้นหาเพลง ด้วยการวัดค่าความคล้ายเชิงมุมกับเวกเตอร์ในตัวแบบ ทำโดยการสุ่มเลือกตัดคำในคำสอบถามออกไป แต่ยังคงลำดับของคำเอาไว้ พบว่าขนาดคำสอบถามที่ร้อยละ 30 - 50 ของเนื้อเพลงทั้งหมด เป็นขนาดที่เหมาะสมที่สุดที่ส่งผลให้ตัวแบบที่ถูกสร้างด้วยเวกเตอร์ขนาด 200 ทำงานอย่างมีประสิทธิภาพ

2. การพัฒนาระบบการค้นหาเพลงจากเสียงร้องด้วยตัวแบบการค้นหาเชิงความหมายจากเทคนิค Doc2Vec สามารถอภิปรายผลได้ดังนี้ จากผลการทดลองเพื่อพัฒนาประสิทธิภาพของตัวแบบฯ ทำให้ผู้วิจัยทราบว่าในการรับข้อมูลเสียงร้องจากผู้สืบค้นผ่านโปรแกรมประยุกต์นั้น ผู้ใช้งานจะต้องร้องเนื้อเพลงไม่น้อยไปกว่าร้อยละ 30 ของเพลงต้นฉบับ และไม่จำเป็นต้องร้องถูกต้องในทุกคำ ซึ่งทำให้ผู้ใช้งานสามารถกำหนดได้ว่าเมื่อต้องการสืบค้นเพลงที่ต้องการจะต้องทำการเนื้อเพลงเป็นความยาวเท่าใดเพื่อให้ผลของการสืบค้นมีความแม่นยำ ซึ่งสอดคล้องกับแนวคิดที่ Guli และ Tiwari [12] ได้นำเสนอในบทสรุปของการวิจัยว่าการทำให้กระบวนการค้นหาเพลงง่ายขึ้น หากผู้ใช้งานสามารถทำเพียงแค่การพูดเนื้อร้องของเพลงที่พวกเขาต้องการ และแบบจำลองที่ถูกพัฒนาไว้จะทำการเลือกเพลงที่คล้ายกันขึ้นมาได้อย่างทันที เพื่อให้ผู้ใช้งานสามารถเข้าถึงเพลงที่ต้องการได้โดยไม่มีความยุ่งยาก

ข้อเสนอแนะ

ชุดข้อมูลที่ใช้ในการวิจัยครั้งนี้ผู้วิจัยเลือกใช้เนื้อเพลงประเภทเพลงลูกทุ่งเพียงอย่างเดียว ไม่ได้รวมเนื้อเพลงประเภทอื่นเข้ามาด้วย ทำให้ข้อมูลมีความสัมพันธ์กันสูงมากเนื่องจากการใช้คำ และจัดวางบริบทแวดล้อมคำใกล้เคียงกันซึ่งพิจารณาจากประสิทธิภาพของตัวแบบที่มีค่าสูง แสดงว่าความสัมพันธ์ของเนื้อเพลงเหล่านี้อาจส่งผลต่อประสิทธิภาพของตัวแบบ ซึ่งการทำงานวิจัยต่อไปอาจเป็นการทดสอบว่าเมื่อใช้เนื้อเพลงที่มีความแตกต่างกันหลายกลุ่มประเภทมาใช้สำหรับสอนตัวแบบ จะส่งผลต่อประสิทธิภาพของตัวแบบอย่างไร การใช้คำสอบถามที่มีความยาวไม่ต่ำกว่าร้อยละ 30 ของเนื้อร้องต้นฉบับเพื่อให้ได้การสืบค้นที่มีความถูกต้องนั้น สำหรับผู้ใช้งานบางท่านอาจไม่อาจยังไม่สะดวกมากนัก ดังนั้นสำหรับการศึกษาวิจัยต่อไปอาจมุ่งประเด็นในเรื่องของการหาทางทำให้ตัวแบบทำงานได้แม่นยำมากขึ้นโดยใช้ขนาดคำสอบถามที่สั้นลงกว่าเดิม

เอกสารอ้างอิง

- [1] สมจิน เปียโคสูง, และนิศาชล จันทน์ศรี. (2553). ระบบนำทางความรู้เพื่อการเข้าถึงเนื้อหาในสื่อสิ่งพิมพ์. *วารสารสารสนเทศศาสตร์*, 28(3), 9-20.
- [2] บุษบงก์ คชินทรโรจน์, เดือนเพ็ญ ธีรวรรณวิวัฒน์ และพาชิตชนิต ศิริพานิช. (2564). การสร้างระบบคัดกรองข้อความการเกลียดกลัวคนต่างชาตินบนทวิตเตอร์ในช่วงการแพร่ระบาดของโรคติดเชื้อไวรัสโคโรนา 2019. *Thai Journal of Operations Research: TJOR*, 9(1), 31-44.
- [3] ภูมิรพี ภูมิคำ. (2562). เทคนิคการเรียนรู้เชิงลึกเพื่อวิเคราะห์ความรู้สึกจากผู้ใช้ผลิตภัณฑ์ (วิทยานิพนธ์ปริญญาโทบริหารธุรกิจ). นครราชสีมา: มหาวิทยาลัยเทคโนโลยีสุรนารี. สืบค้นจาก <http://sutir.sut.ac.th:8080/jspui>

- [4] Mikolov T., Chen K., Corrado G.S., & Dean J. (2013). Efficient Estimation of Word Representations in Vector Space. ICLR.
- [5] Le Q., Mikolov T. (2014). Distributed Representations of Sentences and Document. *Proceedings of the 31st International Conference on Machine Learning*, Beijing, China.
- [6] Dai A.M., Olah C., & Le Q.V. (2015). Document Embedding with Paragraph Vectors. *ArXiv*, abs/1507.07998.
- [7] Lau J.H., & Baldwin T. (2016). An Empirical Evaluation of doc2vec with Practical Insights into Document Embedding Generation. *Proceedings of the 1st Workshop on Representation Learning for NLP (pp. 78–86)*, Berlin, Germany: Association for Computational Linguistics.
- [8] Kaothanthong N., Kongyoung S., & Theeramunkong T. (2021). Headline2Vec: A CNN-based Feature for Thai Clickbait Headlines Classification. *International Scientific Journal of Engineering And Technology*, 5(1), 20-31.
- [9] Alshammeria M., Atwella E., Alsalkaa M.A. (2021). Detecting Semantic-based Similarity Between Verses of The Quran with Doc2vec. *Procedia Computer Science*, 189, 351-358.
- [10] Budiarto A., Rahutomo R., Putra H. N., Cenggoro T. W., Kacamarga M. F., & Pardamean B. (2021). Unsupervised News Topic Modelling with Doc2Vec and Spherical Clustering. *Procedia Computer Science*, 179, 40-46.
- [11] Patra B.G., Das D., and Bandyopadhyay S. (2017). Retrieving Similar Lyrics for Music Recommendation System. *Proceeding of Conference on Natural Language Processing (pp. 290-297)*, Kolkata, India: NLP Association of India (NLP AI).
- [12] Gali N., Tiwari V. (2021). Speech and Lyric-base Doc2Vec Music Recommendation System. *International Journal of Engineering Research & Technology (IJERT)*, Special Issue 9(8), 67-70.