



บทบาทของบทวิพากษ์วิจารณ์จากลูกค้าบนสื่อสังคมต่อการสนับสนุนธุรกิจท่องเที่ยวในประเทศไทย:
เทคนิคการวิเคราะห์ข้อมูลขนาดใหญ่

A Role of Customer Reviews on Social Media to Promoting Travel Businesses in
Thailand: A Technique for Analyzing Big Data

พรภัทร์ ศิริธรรมกุล

งานวิจัยนี้ได้รับทุนสนับสนุนจากงบประมาณเงินรายได้ ประจำปีงบประมาณ พ.ศ.2566

คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร

ชื่อเรื่อง : บทบาทของบทวิพากษ์วิจารณ์จากลูกค้าบนสื่อสังคมต่อการสนับสนุนธุรกิจ
ท่องเที่ยวในประเทศไทย: เทคนิคการวิเคราะห์ข้อมูลขนาดใหญ่

ผู้วิจัย : ดร. พรภัทร์ ศิริธรรมกุล สาขาวิชาวิศวกรรมคอมพิวเตอร์
คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเทคโนโลยี
ราชมงคลพระนคร

พ.ศ. : 2566

บทคัดย่อ

งานวิจัยนี้มุ่งเน้นเพื่อจะหาข้อมูลสำหรับการฟื้นฟูการท่องเที่ยวในประเทศไทยภายหลังสถานการณ์โควิด-19 เนื่องจากประเทศไทยมีรายได้หลักส่วนหนึ่งจากการท่องเที่ยว การปิดประเทศเนื่องจากสถานการณ์โรคระบาดทำให้ประเทศขาดรายได้เป็นจำนวนมหาศาล และส่งผลให้ระบบเศรษฐกิจเกิดการชะลอตัว การทำแผนฟื้นฟูควรต้องอยู่บนพื้นฐานของข้อมูลความเป็นจริงและมีการวิเคราะห์ข้อมูลที่แม่นยำ นักวิจัยจึงมีความคิดที่จะนำข้อมูลที่เกี่ยวข้องกับการท่องเที่ยวมาวิเคราะห์ให้เห็นภาพจุดแข็งและข้อจำกัดของการท่องเที่ยวในประเทศไทยผ่านมุมมองของนักท่องเที่ยวต่างชาติ โดยข้อมูลนี้ถูกนำมาจากบทวิพากษ์วิจารณ์ของนักท่องเที่ยวบนแพลตฟอร์มสื่อสังคมซึ่งมีข้อดีคือ เป็นข้อมูลเปิด (open-source data) ที่ไม่มีค่าใช้จ่าย และมีจำนวนที่มากพอที่จะทำให้ผลของการวิเคราะห์ข้อมูลได้สะท้อนความเป็นจริงเกี่ยวกับความพึงพอใจของนักท่องเที่ยวในแง่มุมต่าง ๆ

ในแง่ของเทคนิคที่ใช้เพื่อการวิเคราะห์ข้อมูล งานวิจัยนี้ใช้ Retrieval-Augmented Generation (RAG) เพื่อเพิ่มความสามารถของแบบจำลองภาษาขนาดใหญ่ (Large Language Models: LLMs) โดย RAG นี้สร้างขึ้นจาก Llama2 ซึ่งเป็น LLM ที่พัฒนาโดย Meta AI และข้อมูลที่น่ามาวิเคราะห์นี้มาจากบทวิพากษ์วิจารณ์ของนักท่องเที่ยวบน TripAdvisor ซึ่งเป็นเว็บไซต์เกี่ยวกับการท่องเที่ยว งานวิจัยนี้จะแสดงให้เห็นถึงประสิทธิภาพของ RAG ซึ่งมีประสิทธิภาพดีกว่า LLMs ได้แก่ ChatGPT, Perplexity, Gemini และ Copilot ในแง่ของการให้ผลการวิเคราะห์เป็นข้อมูลเชิงลึกที่ LLMs ดังกล่าวไม่เคยวิเคราะห์ได้มาก่อน

ท้ายที่สุด งานวิจัยนี้ใช้ผลลัพธ์จาก RAG มาใช้เพื่อหาแนวทางการแก้ปัญหาการท่องเที่ยว รวมถึงได้แนะนำแนวทางการปฏิบัติที่ใช้ในแต่ละสถานที่เพื่อส่งเสริมการท่องเที่ยวต่อไป

Title : A Role of Customer Reviews on Social Media to Promoting Travel Businesses in Thailand: A Technique for Analyzing Big Data

Researcher : Dr. Pornpat Sirithumgul Department of Computer Engineering,
Faculty of Engineering, Rajamangala
University of Technology Phra Nakhon

Year : 2023

ABSTRACT

This research aims to gather information for tourism recovery in Thailand after the COVID-19 pandemic. As tourism is a major source of revenue for Thailand, the country-wide lockdown due to the pandemic has resulted in substantial losses and economic slowdown. Tourism recovery plans need to be based on real-world data and accurate analysis. Therefore, this study utilizes open-source data from tourist reviews on social media platforms to analyze the strengths and limitations of Thai tourism from the perspective of international tourists. This dataset is abundant and cost-effective, allowing for reliable analysis of tourist satisfaction across various aspects.

Regarding the text analysis technique, this study employs Retrieval-Augmented Generation (RAG) to enhance the capabilities of large language models (LLMs). This RAG is built upon Llama2, an LLM developed by Meta AI, and analyzes tourist reviews from TripAdvisor, an online travel website. This research study demonstrates RAG's superior performance compared to LLMs, including ChatGPT, Perplexity, Gemini, and Copilot in identifying insights previously unreported by these models.

Finally, this study applies the findings from RAG to pinpoint tourism recovery strategies and recommends site-specific practices for promoting tourism in the future.

กิตติกรรมประกาศ

งานวิจัยนี้สำเร็จด้วยทุนวิจัยจากงบประมาณรายได้ประจำปีงบประมาณ 2566 [รหัสทุนวิจัย 66 – 106 – 02/3] มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร อีกทั้ง ผู้วิจัยขอขอบพระคุณคณาจารย์จากภาควิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัย คณาจารย์จากสาขาวิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร และคณาจารย์จากคณะวิทยาศาสตร์และเทคโนโลยี แขนงวิชาเทคโนโลยีสารสนเทศ และการสื่อสาร มหาวิทยาลัยสุโขทัยธรรมาธิราช ที่เอื้อเฟื้อคำแนะนำ ข้อมูล และสถานที่ที่ใช้ในงานวิจัย มา ณ โอกาสนี้

พรภัทร์ ศิริธรรมกุล



สารบัญ

	หน้า
บทคัดย่อ	II
ABSTRACT	III
กิตติกรรมประกาศ	IV
สารบัญ	V
สารบัญตาราง	VII
สารบัญภาพ	VIII
บทที่ 1	9
บทนำ	9
1.1 ความเป็นมาและความสำคัญของงานวิจัย	9
1.2 วัตถุประสงค์ของการวิจัย	11
1.3 การทบทวนวรรณกรรม/สารสนเทศที่เกี่ยวข้อง	11
1.4 ขอบเขตการวิจัย	13
1.5 ประโยชน์ที่ได้รับจากงานวิจัย	13
บทที่ 2	14
ทฤษฎีที่เกี่ยวข้อง	14
บทที่ 3	19
โมเดลสำหรับการประมวลผลข้อมูล	19
บทที่ 4	22
การทดลอง และข้อมูลที่ใช้ในการทดลอง	22
บทที่ 5	27
ผลการทดลอง	27

บทที่ 6	33
บทวิเคราะห์และงานวิจัยสืบเนื่องในอนาคต	33
บทที่ 7	41
บทสรุปงานวิจัย	41
บรรณานุกรม	42
ประวัติผู้วิจัย	44



สารบัญตาราง

ตารางที่	หน้า
4.1 หมวดหมู่คำวิพากษ์วิจารณ์จาก TripAdvisor	23
4.2 หมวดหมู่อารมณ์ความรู้สึกของบทวิพากษ์วิจารณ์จาก TripAdvisor	24
5.1 เอาต์พุต Response จาก RAG และ LLMs	28
5.2 ค่า Reliability ของ RAG และ LLMs	30
5.3 ค่า Reliability จำแนกตามบทวิพากษ์วิจารณ์ของแต่ละสถานที่ท่องเที่ยว	31
6.1 แสดงแนวคิดในการพัฒนาสำหรับสถานที่ท่องเที่ยวทั้ง 6 แห่งในกรุงเทพฯ	38



สารบัญภาพ

ภาพที่	หน้า
2.1 โค้ดภาษา Python เพื่อเรียกใช้งาน Vader ในการประมวลผล	15
2.2 ผลลัพธ์การประมวลผลโดยใช้ Vader	15
2.3 แบบจำลอง semantic feature space ในระนาบ 2 มิติ	16
2.4 แบบจำลอง semantic feature space เมื่อมีการเพิ่มจำนวนคำในฐานข้อมูล	17
2.5 ตัวอย่างการคำนวณ embeddings และ cosine similarity โดยใช้ spaCy	18
2.6 ผลลัพธ์การคำนวณ cosine similarity ระหว่างสองคำที่กำหนด	18
3.1 ขั้นตอนการสร้าง Vector Database	19
3.2 ขั้นตอนการสร้างคำตอบจาก RAG	20
3.3 ตัวอย่างการสร้าง question prompt เพื่อป้อนให้กับ Llama2	21
4.1 กระบวนการทดลองของงานวิจัยในภาพรวม	25



บทที่ 1

บทนำ

บทนี้กล่าวถึงความเป็นมาและความสำคัญของงานวิจัย วัตถุประสงค์ ขอบเขตงานวิจัย และประโยชน์ที่ได้รับจากงานวิจัย ซึ่งมีรายละเอียดดังต่อไปนี้

1.1 ความเป็นมาและความสำคัญของงานวิจัย

งานวิจัยนี้มุ่งหวังที่จะใช้ความสามารถของ Large Language Models (LLMs) ในสองด้านหลักคือ การหาข้อมูลโดยสรุปของเอกสาร (document summarization) และการวิเคราะห์ภาษา (text analytics) ในการพิจารณาคำตอบของคำถาม (question prompt) ที่ป้อนให้กับ LLMs คือ “อภิปรายข้อดีและข้อเสียของการเยี่ยมชม [สถานที่ท่องเที่ยว] ในกรุงเทพฯ ประเทศไทย” โดย LLMs ที่นำมาพิจารณาในการทดลองของงานวิจัยนี้ได้แก่ ChatGPT, Perplexity, Gemini และ Copilot โดย LLMs เหล่านี้จะสร้างคำตอบบนพื้นฐานของข้อมูลมหาศาลในขั้นตอนการสอนโมเดล (training data) อย่างไรก็ตามคำตอบที่ได้จาก LLMs บางคำตอบนั้นขาดความน่าเชื่อถือ เพราะคำตอบเหล่านี้มีข้อผิดพลาดที่เรียกว่า ‘ภาพหลอน’ (hallucination) และเป็นคำตอบที่ล้าสมัย (outdated response)

การเกิด hallucination หมายถึงการสร้างคำตอบของ LLMs ที่ไม่มีความแม่นยำ ไม่เป็นความจริง และไม่มีความเกี่ยวข้องเชื่อมโยงกับคำถามนำ (question prompt) ที่ป้อนให้กับ LLMs ปัญหาลักษณะนี้เกิดขึ้นจากข้อมูลที่ใช้สอนโมเดล LLMs มีลักษณะจำเพาะมากเกินไป (overfitting) หรือข้อมูลที่ใช้ในการสอนโมเดลมีความขัดแย้งกันเอง (conflicting data) ในทำนองเดียวกัน คำตอบที่ล้าสมัย (outdated response) จาก LLMs เกิดจากข้อมูลที่ใช้สอนโมเดล LLMs มีความล้าสมัย ปัญหาดังที่กล่าวข้างต้นส่งผลกระทบต่อความน่าเชื่อถือของคำตอบของ LLMs ทั้งที่ประสิทธิภาพในการประมวลผลข้อมูลของ LLMs เหล่านี้ยังมีความชาญฉลาด

เพื่อเป็นการแก้ปัญหาของ LLMs ที่กล่าวมาข้างต้น งานวิจัยนี้จึงเสนอการออกแบบโมเดล Retrieval-Augmented Generation (RAG) ที่ทำงานบนพื้นฐานของ LLM ชื่อ Llama ที่ออกแบบโดย Meta AI

อนึ่ง RAG ที่ถูกพัฒนาขึ้นจากงานวิจัยนี้ ได้ถูกนำมาใช้ประมวลผลข้อมูลจาก TripAdvisor โดยเอาต์พุตจาก RAG สามารถแสดงให้เห็นข้อมูลจากการประมวลผล 3 ลักษณะดังนี้

- 1) **การค้นพบความรู้ใหม่ที่ไม่เคยค้นพบได้จาก LLMs** ในที่นี้อ้างอิงจากเอาต์พุตที่ได้จาก RAG ซึ่งต่างจากเอาต์พุตที่ได้จาก ChatGPT, Perplexity, Gemini และ Copilot
- 2) **เอาต์พุตที่ได้จาก RAG จะถูกใช้เพื่อยืนยันข้อมูลที่ถูกต้องจาก LLMs** ทั้ง 4 ตัวคือ ChatGPT, Perplexity, Gemini และ Copilot โดยสังเกตจากเอาต์พุตจาก RAG และ LLMs ที่สอดคล้องกัน
- 3) **เอาต์พุตที่ได้จาก LLM ตัวใดตัวหนึ่งมีแนวโน้มจะเป็นข้อผิดพลาด** โดยสังเกตจาก เอาต์พุตจาก LLM ตัวนั้นที่ไม่สอดคล้องกับเอาต์พุตจาก LLMs ตัวอื่น และเอาต์พุตจาก RAG

งานวิจัยนี้จึงมุ่งหวังที่จะนำเสนอและสร้างโมเดล RAG และทำการเปรียบเทียบประสิทธิภาพการทำงานกับ LLMs ที่มีชื่อเสียงทั้ง 4 ตัว ได้แก่ ChatGPT, Perplexity, Gemini และ Copilot โดยการสร้าง RAG จะใช้หลักการของการวิเคราะห์ความรู้สึก (sentiment analysis) ซึ่งจะแปลงข้อมูลตัวอักษร (text) เป็นข้อมูลตัวเลข (numerical embeddings) และจัดเก็บในฐานข้อมูลที่เรียกว่า “vector database” โดยข้อมูล embeddings นี้จะถูกใช้เป็นดัชนี (index) สำหรับการสืบค้นข้อมูลจากฐานข้อมูลออกมาใช้ ในช่วง run time ข้อมูล embeddings จะถูกสืบค้น (retrieve) จาก vector database และถูกแปลงเป็นตัวอักษร ก่อนจะถูกนำไปใช้เป็นส่วนหนึ่งของการสร้าง “คำถามนำ” (question prompt) ก่อนที่จะถูกป้อนให้กับ LLM ที่เลือกใช้ – ในที่นี้คือ Llama 2 ที่พัฒนาโดย Meta AI เพื่อทำการประมวลผลข้อมูล และคืนค่าเอาต์พุตที่ได้จากการประมวลผล

เอาต์พุตที่ได้จาก RAG จะถูกนำไปพิจารณาความน่าเชื่อถือ โดยเปรียบเทียบกับเอาต์พุตที่ได้จาก LLMs – ChatGPT, Perplexity, Gemini และ Copilot ในขั้นตอนสุดท้ายของการวิจัย เอาต์พุตส่วนที่พิสูจน์ความน่าเชื่อถือแล้ว ได้ถูกนำไปเป็นข้อมูลสำคัญสำหรับการปรับปรุงนโยบายการพัฒนาการท่องเที่ยวให้กับสถานที่ท่องเที่ยวที่สำคัญในประเทศไทย

1.2 วัตถุประสงค์ของการวิจัย

- 1) เพื่อสร้างโมเดลปัญญาประดิษฐ์สำหรับประมวลผลภาษาธรรมชาติ
- 2) เพื่อหาบทสรุปความพึงพอใจ และปัจจัยที่ทำให้นักท่องเที่ยวชื่นชอบ และถูกมองว่าเป็นอุปสรรคของการท่องเที่ยวสถานที่ท่องเที่ยวยอดนิยมในประเทศไทย
- 3) เพื่อเสนอบทสรุปจากการวิเคราะห์ข้อมูลที่ได้จากสื่อสังคมสำหรับทำแผนฟื้นฟูการท่องเที่ยวในประเทศไทย

1.3 การทบทวนวรรณกรรม/สารสนเทศที่เกี่ยวข้อง

ในหัวข้อนี้จะกล่าวถึงหลักการสร้าง Retrieval-Augmented Generation (RAG) และตัวอย่างการประยุกต์ใช้ RAG ในการวิจัยก่อนหน้า ได้แก่ การประยุกต์ใช้ RAG ในวงการสาธารณสุข การเงิน การศึกษา และอีคอมเมิร์ซ (e-commerce) ตามลำดับ

Retrieval-Augmented Generation (RAG) เป็นโมเดลประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) ที่ถูกนำมาใช้เพื่อลดข้อจำกัดของ Large Language Models (LLMs) ในการวิเคราะห์ข้อมูล และสร้างข้อความตัวอักษรเพื่ออธิบายความรู้ในโดเมน (knowledge domain) ที่ต้องการ ภายในโครงสร้างของ RAG มี LLM เป็นส่วนประกอบสำหรับประมวลผลข้อมูล อย่างไรก็ตาม RAG มิได้ใช้ข้อมูลที่ใช้ในการสอน (training data) สำหรับโมเดล LLM หากแต่ RAG จะใช้ข้อมูลชุดใหม่ที่ได้รับการแก้ไขให้เป็นปัจจุบันที่สุดในชุดโดเมนความรู้ นั้น ๆ

โมเดล RAG แบ่งออกเป็น 2 ส่วนหลัก เรียกว่า Retriever และ Generator โดยส่วนแรกจะทำการวิเคราะห์คำถามจากผู้ใช้งาน (user's query) และทำการสืบค้น (retrieve) ข้อมูลออกมาจากฐานข้อมูล ส่วนที่เรียกว่า Generator ของ RAG ทำหน้าที่สร้าง (generate) ข้อความที่ถูกต้องที่ใช้เป็นคำตอบของคำถาม (user's query) จากผู้ใช้งาน ในส่วนของ Generator นี้จะใช้ LLM ในการประมวลผล โดยการป้อนคำถามนำ (question prompt) ที่เกิดจากการรวมตัวกันระหว่างคำถามจากผู้ใช้งาน (user's query) และข้อมูลจากฐานข้อมูลที่ถูกสืบค้นออกมาโดย Retriever

RAG สามารถทำงานได้ในหน้าที่เดียวกับ LLMs ไม่ว่าจะเป็นหน้าที่ของการตอบคำถาม (question-answering) การหาบทสรุปข้อมูล (summarizing data) การสร้างบทสนทนาโต้ตอบกับผู้ใช้งาน (generating dialogues interacting with users) อย่างไรก็ตาม ปัญหาที่จะเกิดกับ LLMs โดยหลักแล้วมีอยู่ 2 ประการ ได้แก่ การสร้างคำตอบที่ล้าสมัย (outdated response) และการสร้างคำตอบที่ไม่เกี่ยวข้องกับคำถามเรียกว่า ‘ภาพหลอน’ (hallucination) และการสร้างคำตอบที่ล้าสมัย (outdated response) ซึ่งปัญหาทั้งสองประการนี้จะลดลงอย่างมากเมื่อใช้ RAG เนื่องจากการประมวลผลของ RAG จะใช้ข้อมูลที่เป็นปัจจุบันมากที่สุดและใช้ข้อมูลที่อยู่ในบริบท (context-aware data) ในการประมวลผล

ในอีกแง่มุมหนึ่ง RAG ยังใช้ทรัพยากรด้านงบประมาณ และราคาที่ต้องจ่ายน้อยกว่า LLMs เนื่องจาก LLMs ต้องการเวลาในกระบวนการสอน (model-training phase) เพื่อเตรียมความพร้อมของโมเดล และต้องการข้อมูลสำหรับการสอน (training data) จำนวนมากในการสร้างโมเดล ในทางตรงกันข้าม RAG ต้องการเวลาและงบประมาณน้อยกว่ามาก เนื่องจาก RAG จะใช้การประมวลผลแบบ real-time โดยการสืบค้นข้อมูลจากแหล่งข้อมูลซึ่งเป็นประโยชน์กับการประมวลผลและผนวกข้อมูลเข้ากับคำถามจากผู้ใช้งาน (user's query) จากนั้นจะทำการป้อนข้อมูลที่ผนวกกันแล้วทั้งกลุ่มป้อนให้กับ LLM ที่ฝังตัวใน RAG (embedded LLM) ที่ RAG เลือกใช้ในการประมวลผล

การพัฒนา RAG เกิดขึ้นในหลากหลายอุตสาหกรรม อาทิ สาธารณสุข การเงิน การศึกษา และอีคอมเมิร์ซ ในงานวิจัยของ Ghadban และคณะ (ปี 2023) สร้าง RAG ชื่อ “*SMARThealthGPT*” สำหรับการตอบคำถามให้กับกลุ่มบุคลากรทางการแพทย์ในช่วงเวลาการฝึกอบรมด้านการดูแลสุขภาพสตรีมีครรภ์ (maternal care)

อุตสาหกรรมการเงิน เป็นอีกหนึ่งวงการที่ใช้เครื่องมือถาม-ตอบ (question-answering tool) เพื่อตอบคำถามเฉพาะวงการการเงินที่ประกอบด้วยคำศัพท์เฉพาะทาง (financial jargon) ในงานวิจัยของ Zhang และคณะ (ปี 2023) สร้างเครื่องมือชื่อ “*FinBERT*” สำหรับการวิเคราะห์บทความทางการเงิน ซึ่งให้ผลลัพธ์เป็นบทความที่อ่านเข้าใจง่ายโดยคนนอกวงการการเงิน

ทั้ง LLMs และ RAG มีบทบาทสำคัญอย่างมากในการพัฒนา AI เพื่อการศึกษา งานวิจัยของ Dong (ปี 2023) เป็นตัวอย่างหนึ่งของการสร้าง AI เพื่อเป็นผู้ช่วยครู (tutor) สอนพิเศษ เนื่องจาก AI มีการใช้ภาษาธรรมชาติเพื่อโต้ตอบคำถามจากนักเรียน โดยคำตอบจาก AI จะอ้างอิงจากเอกสารประกอบการเรียนที่ใช้ในชั้นเรียน

RAG มีบทบาทสำคัญในการส่งเสริมการให้บริการลูกค้าของอีคอมเมิร์ซ ในงานวิจัยของ Kulkarni และคณะ (ปี 2024) ใช้ RAG ในการพัฒนาแชทบอท (chatbot) ให้มีประสิทธิภาพมากขึ้น เพื่อให้แชทบอทตอบคำถามลูกค้าเกี่ยวกับบัตรเครดิต งานวิจัยนี้แสดงให้เห็นถึงค่าใช้จ่ายที่มีราคาต่ำแต่ให้คำตอบที่ถูกต้องแม่นยำมากขึ้น

งานวิจัยฉบับนี้มุ่งหวังที่จะสร้าง RAG ในการวิเคราะห์ข้อมูลการท่องเที่ยว เพื่อแสดงให้เห็นประสิทธิภาพของ RAG ในชุดข้อมูลที่แตกต่างจากงานวิจัยก่อนหน้านี้ และแสดงให้เห็นถึงการนำผลลัพธ์จาก RAG เพื่อนำไปใช้ประโยชน์ในอุตสาหกรรมการท่องเที่ยวต่อไป

1.4 ขอบเขตการวิจัย

- 1) ออกแบบและสร้างโมเดล RAG สำหรับประมวลผลบทวิพากษ์วิจารณ์เกี่ยวกับการท่องเที่ยว
- 2) ทดสอบประสิทธิภาพและความน่าเชื่อถือของโมเดลที่สร้างขึ้น เพื่อศึกษาแนวโน้มของการนำผลลัพธ์ที่ได้จากโมเดลไปใช้งานต่อ
- 3) นำผลลัพธ์ที่ได้จากการประมวลผลโดยโมเดลที่สร้างไปใช้เพื่อสร้างนโยบายส่งเสริมการท่องเที่ยว

1.5 ประโยชน์ที่ได้รับจากงานวิจัย

โมเดล retrieval-augmented generation (RAG) สำหรับประมวลผลข้อมูล และบทวิเคราะห์ข้อดีและข้อจำกัดของสถานที่ท่องเที่ยวที่สำคัญในประเทศไทย และแนวทางการปรับปรุงสถานที่ท่องเที่ยวเพื่อสร้างประสบการณ์ที่ดีให้กับนักท่องเที่ยว



บทที่ 2

ทฤษฎีที่เกี่ยวข้อง

งานวิจัยนี้ใช้เทคนิคสำคัญ ได้แก่ การวิเคราะห์อารมณ์ความรู้สึก (sentiment analysis) ที่ถ่ายทอดในบทความภาษาอังกฤษ และฐานข้อมูลเวกเตอร์ (vector database) โดยเทคนิค sentiment analysis เป็นส่วนหนึ่งของการวิเคราะห์ภาษาธรรมชาติ (Natural Language Processing: NLP) ที่ใช้โปรแกรมภาษา ไพทอน (Python) เพื่อการวิเคราะห์คำศัพท์ที่บ่งบอกอารมณ์ความรู้สึก – พอใจ หรือไม่พอใจ เมื่อกำหนดหัวข้อเฉพาะ ซึ่งในบริบทของงานวิจัยนี้คือความรู้สึกของนักท่องเที่ยวที่ถ่ายทอดในสื่อสังคมเพื่ออธิบายความพึงพอใจเกี่ยวกับสถานที่ท่องเที่ยวในกรุงเทพมหานคร สำหรับ vector database เป็นเทคนิคที่ใช้เก็บข้อมูล โดยมีการแปลงข้อมูลตัวอักษร (text) เป็นข้อมูลตัวเลข (numerical data) ที่อธิบายมิติสำคัญบางประการของข้อมูลก่อนการจัดเก็บ ข้อมูลตัวเลขเหล่านี้จะถูกใช้เป็นดัชนีในการจัดเก็บข้อมูลที่มีความคล้ายคลึงกัน และทำให้การสืบค้นข้อมูลที่มีคุณลักษณะใกล้เคียงกันสะดวก และรวดเร็วยิ่งขึ้น รายละเอียดของการประมวลผล sentiment analysis และหลักการออกแบบ vector database มีรายละเอียด ดังนี้

1. การวิเคราะห์อารมณ์ความรู้สึก (sentiment analysis) เป็นส่วนหนึ่งของการประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP) โดยผลลัพธ์ของการวิเคราะห์ คือ บทสรุปว่าบทความร้อยแก้วที่สนใจแสดงอารมณ์ของผู้สื่อสารในหมวดหมู่ใดระหว่าง บวก (positive) กลางๆ (neutral) หรือลบ (negative) โดยบทวิเคราะห์ sentiment analysis นิยมใช้เพื่อวิเคราะห์ความพึงพอใจของกลุ่มคนบนสื่อออนไลน์ (social media posts) และบทวิพากษ์วิจารณ์ของลูกค้าที่ซื้อสินค้าออนไลน์ (customer reviews) เป็นต้น

เครื่องมือที่ใช้สำหรับการวิเคราะห์ sentiment analysis ในงานวิจัยนี้คือ Vader (Valence Aware Dictionary and sEntiment Reasoner) ซึ่งเป็นโมเดลวิเคราะห์ประเภท rule-based ที่ยึดตามคำศัพท์ในพจนานุกรม (dictionary of words) ที่ถูกใช้สำหรับฝึกสอนโมเดล (pre-trained model) เอาต์พุตจากโมเดลนี้จะเป็นคะแนนที่ได้จากการประเมิน โดยค่าจะอยู่ระหว่าง +4 ถึง -4 โดย +4 แสดงถึงคำศัพท์ในบทความมีแนวคิดเป็นบวก (positive) มากที่สุด และ -4 แสดงว่าคำศัพท์ที่ใช้ในการสร้างบทความแสดงแนวคิดที่เป็นลบ (negative) มากที่สุด

งานวิจัยนี้ใช้การโปรแกรมภาษา Python ในการประมวลผล sentiment analysis และใช้ Natural Language Toolkit (NLTK) library ซึ่งได้รวมความสามารถในการประมวลผลของ Vader ไว้ด้วยแล้วในการสร้างซอฟต์แวร์สำหรับประมวลผลข้อมูล ทั้งนี้ Python Class ชื่อ “SentimentIntensityAnalyzer” จาก module “nltk.sentiment.vader” จะถูกเรียกใช้ใน

โค้ดภาษา Python และมีการเรียกใช้เมธอด `polarity_scores()` ที่มีการ return ค่าเอาต์พุตเป็นชุด dictionary ที่เก็บค่าตัวแปร 4 ค่า ได้แก่ `neg`, `neu`, `pos` และ `compound` ดังนี้

- ค่าของตัวแปร `neg` เก็บค่าคะแนน sentiment score ที่เป็นลบ มีค่าระหว่าง 0 ถึง 1
- ค่าของตัวแปร `neu` เก็บค่าคะแนน sentiment score ที่เป็นกลาง มีค่าระหว่าง 0 ถึง 1
- ค่าของตัวแปร `pos` เก็บค่าคะแนน sentiment score ที่เป็นบวก มีค่าระหว่าง 0 ถึง 1
- ค่าของตัวแปร `compound` เก็บค่าคะแนน sentiment score ในภาพรวม มีค่าระหว่าง -1 ถึง 1

ตัวอย่างต่อไปนี้จะแสดงโค้ดภาษา Python ¹ ที่เรียกใช้ Vader เป็นเครื่องมือในการประมวลผลประโยคตัวอย่าง 3 ประโยค ได้แก่ “I love this product! It works great and is very affordable.”, “This product is okay. It gets the job done, but could be better.” และ “I hate this product. It doesn’t work at all and is a waste of money.” เอาต์พุตจากการประมวลผลทั้ง 3 ประโยคโดยใช้ Vader ปรากฏว่าประโยคแรกมีค่า sentiment เป็นบวก, ประโยคที่สองมีค่า sentiment เป็นกลางค่อนข้างไปทางบวก และประโยคที่สามมีค่า sentiment เป็นลบตามลำดับ

```
1 from nltk.sentiment.vader import SentimentIntensityAnalyzer
2
3 analyzer = SentimentIntensityAnalyzer()
4 texts = ["I love this product! It works great and is very affordable.",
5         "This product is okay. It gets the job done, but could be better",
6         "I hate this product. It doesn't work at all and is a waste of money."]
7 for text in texts:
8     scores = analyzer.polarity_scores(text)
9     print(text)
10    print(scores)
```

ภาพที่ 2.1: โค้ดภาษา Python เพื่อเรียกใช้งาน Vader ในการประมวลผล

¹ ตัวอย่างโค้ดภาษา Python ดัดแปลงจากบทความออนไลน์ Vader: A Comprehensive Guide to Sentiment Analysis in Python. <https://medium.com/@rslavanyageetha/vader-a-comprehensive-guide-to-sentiment-analysis-in-python-c4f1868b0d2e#:~:text=Sentiment%20Analysis%20using%20Vader%20in%20Python&text=Sentiment%20analysis%20is%20a%20popular,trained%20model%20for%20sentiment%20analysis.> วันที่ 25 มิถุนายน 2567

```

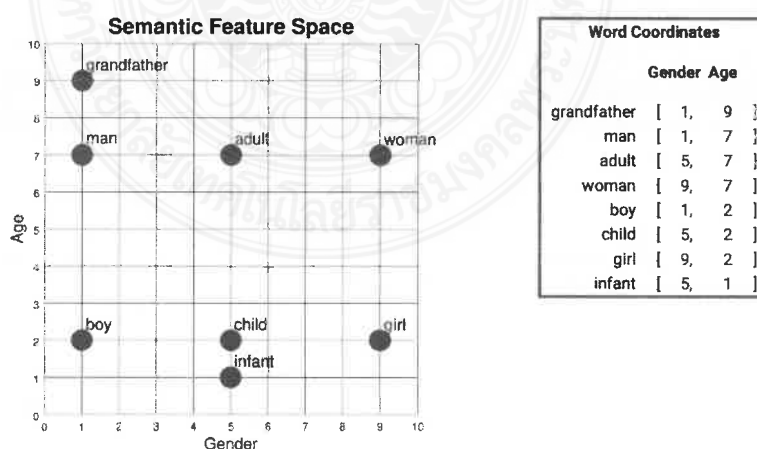
/Users/pornpatsirithumgul/PycharmProjects/sentimentSample/.venv/bin/python
I love this product! It works great and is very affordable.
{'neg': 0.0, 'neu': 0.482, 'pos': 0.518, 'compound': 0.8622}
This product is okay. It gets the job done, but could be better
{'neg': 0.0, 'neu': 0.675, 'pos': 0.325, 'compound': 0.6486}
I hate this product. It doesn't work at all and is a waste of money.
{'neg': 0.371, 'neu': 0.629, 'pos': 0.0, 'compound': -0.7579}

```

ภาพที่ 2.2: ผลลัพธ์การประมวลผลโดยใช้ Vader

- การออกแบบเวกเตอร์ดาต้าเบส (Vector Database) ประกอบด้วย 2 ขั้นตอนหลัก ได้แก่ ขั้นตอนการสร้าง word embedding ซึ่งเป็นการถอดรหัสความหมาย (semantic meaning) ของคำในมิติที่หลากหลาย และแปลงรหัสดังกล่าวเป็นตัวเลขที่เก็บเป็นชุดของข้อมูลเรียกว่า เวกเตอร์ (vector) และขั้นตอนการคำนวณความคล้ายกันของเวกเตอร์ (vector similarity) เนื่องจาก ข้อมูลที่มีความคล้ายกันจะถูกจับกลุ่มให้อยู่ด้วยกันในฐานข้อมูล และถูกดึงออกจากฐานข้อมูลทั้งกลุ่มเพื่อประมวลผล การคำนวณ vector similarity ที่นิยม เช่น cosine similarity

ขั้นตอนการสร้าง word embedding เป็นการถอดรหัสความหมายของคำในมิติต่าง ๆ ให้เป็นตัวเลข ทำให้เกิดการเปรียบเทียบความหมายของคำทำได้โดยการคำนวณ ดังแสดงตัวอย่างในภาพที่ 2.3 ที่แสดงมิติของความหมายของคำใน 2 มิติ ได้แก่ เพศ และอายุ คำเหล่านี้ได้แก่ grandfather, man, boy, adult, child, infant, woman และ girl จะเห็นจาก semantic feature space (ภาพที่ 2.3) ได้ว่า child และ infant มีความใกล้เคียงกันในมิติของอายุ และ girl กับ woman มีความใกล้เคียงกันมากกว่า girl กับ man ในมิติของเพศ

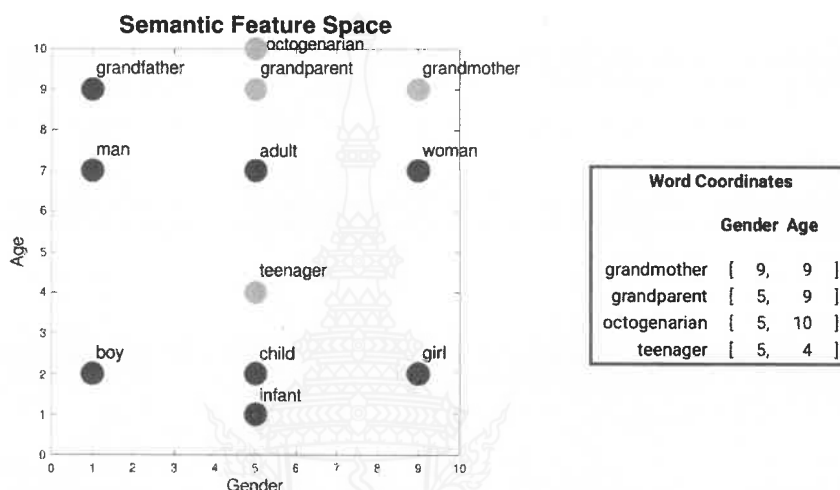


ภาพที่ 2.3: แบบจำลอง semantic feature space ในระนาบ 2 มิติ

ที่มา: Word Embedding Demo: Tutorial.

<https://www.cs.cmu.edu/~dst/WordEmbeddingDemo/tutorial.html> วันที่ 27 มิถุนายน 2567

ในทำนองเดียวกัน การจัดเก็บข้อมูลในฐานข้อมูล และการดึงข้อมูลจากฐานข้อมูลจะพิจารณาจากความใกล้เคียงของคำในทุกมิติที่พิจารณา สมมติต้องการเพิ่มคำต่อไปนี้ลงในระบบ คำเหล่านี้จะถูกนำไปคำนวณจากมิติด้านอายุ และเพศ เพื่อเลือกตำแหน่งที่ใกล้เคียงกับคำศัพท์ที่มีอยู่เดิม เช่นหากต้องการเพิ่มคำว่า octogenarian, grandparent, grandmother และ teenager ตำแหน่งของคำเหล่านี้จะเป็นตามภาพที่ 2.4



ภาพที่ 2.4: แบบจำลอง semantic feature space เมื่อมีการเพิ่มจำนวนคำในฐานข้อมูล

ที่มา: Word Embedding Demo: Tutorial.

<https://www.cs.cmu.edu/~dst/WordEmbeddingDemo/tutorial.html> วันที่ 27 มิถุนายน 2567

ขั้นตอนการคำนวณ **vector similarity** คือการคำนวณ dot product ระหว่างตัวเลขใน vector 2 ชุด เพื่อดูความคล้ายคลึงของข้อมูลใน vector ในการคำนวณที่ใช้ cosine similarity คือการคำนวณค่า $\cos(\theta)$ ระหว่าง vector U และ vector V ดังแสดงในสมการ

$$S_c(U, V) = \cos(\theta) = \frac{U \cdot V}{\|U\| \|V\|}$$

ค่า cosine similarity ที่คำนวณได้จะอยู่ระหว่าง -1 และ 1 โดยการตีความของค่าที่คำนวณได้ เป็นดังนี้

- 1) หากค่าที่คำนวณได้เท่ากับ 1 หมายถึงมุม θ ระหว่างเวกเตอร์ U และ V มีคี่กรเท่ากับ 0 องศา สามารถตีความได้ว่าข้อมูลของทั้งสองเวกเตอร์คล้ายกันมากที่สุด
- 2) หากค่าที่คำนวณได้เท่ากับ 0 หมายถึงมุม θ ระหว่างเวกเตอร์ U และ V มีคี่กรเท่ากับ 90 องศา สามารถตีความได้ว่าข้อมูลของทั้งสองเวกเตอร์ไม่เกี่ยวข้องกัน

- 3) หากค่าที่คำนวณได้เท่ากับ -1 หมายถึงมุม θ ระหว่างเวกเตอร์ U และ V มีดีกรีเท่ากับ 180 องศา สามารถตีความได้ว่าข้อมูลของทั้งสองเวกเตอร์ไม่มีความคล้ายกันเลย

ตัวอย่างต่อไปนี้² เป็นการใช้ spaCy ซึ่งเป็น library ภาษา Python สำหรับประมวลผลภาษาธรรมชาติ (NLP library) เพื่อแปลงคำศัพท์เป็น embedding จากตัวอย่างในโค้ด คำศัพท์เหล่านี้ ได้แก่ dog, cat, apple, tasty, delicious และ truck และคำนวณ cosine similarity ระหว่างคำศัพท์ ผลลัพธ์แสดงให้เห็นคำศัพท์ที่มีความคล้ายกันมาก ได้แก่ ระหว่าง dog กับ cat และ delicious กับ tasty ส่วนคำศัพท์ที่มีความเกี่ยวข้องแต่ไม่คล้ายกัน คือ apple กับ delicious คำศัพท์ที่ไม่มีความคล้าย ได้แก่ dog กับ apple และ truck กับ delicious

```

1 import spacy
2 import numpy as np
3
4 nlp = spacy.load("en_core_web_md")
5 usage:
6 def compute_cosine_similarity(u: np.ndarray, v: np.ndarray) -> float:
7     return (u@v) / (np.linalg.norm(u) * np.linalg.norm(v))
8
9 dog_embedding = nlp.vocab["dog"].vector
10 cat_embedding = nlp.vocab["cat"].vector
11 apple_embedding = nlp.vocab["apple"].vector
12 tasty_embedding = nlp.vocab["tasty"].vector
13 delicious_embedding = nlp.vocab["delicious"].vector
14 truck_embedding = nlp.vocab["truck"].vector
15
16 print("similarity between dog and cat = ", compute_cosine_similarity(dog_embedding, cat_embedding))
17 print("similarity between delicious and tasty = ", compute_cosine_similarity(delicious_embedding, tasty_embedding))
18 print("similarity between apple and delicious = ", compute_cosine_similarity(apple_embedding, delicious_embedding))
19 print("similarity between dog and apple = ", compute_cosine_similarity(dog_embedding, apple_embedding))
20 print("similarity between truck and delicious = ", compute_cosine_similarity(truck_embedding, delicious_embedding))

```

ภาพที่ 2.5: ตัวอย่างการคำนวณ embeddings และ cosine similarity โดยใช้ spaCy

```

/Users/pornpatsirithumgul/PycharmProjects/cosineSimilarity
similarity between dog and cat = 0.8220817
similarity between delicious and tasty = 0.84820914
similarity between apple and delicious = 0.5347654
similarity between dog and apple = 0.22881007
similarity between truck and delicious = 0.0897876

```

ภาพที่ 2.6: ผลลัพธ์การคำนวณ cosine similarity ระหว่างสองคำที่กำหนด

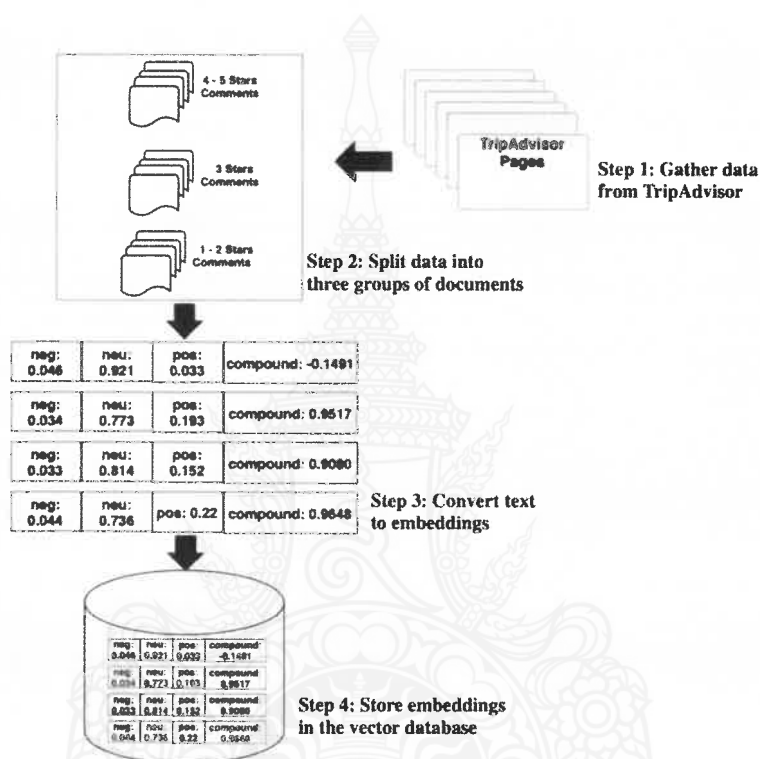
² ตัวอย่างโค้ดภาษา Python ดัดแปลงจากบทความออนไลน์ Embeddings and Vector Databases With ChromaDB.

<https://realpython.com/chromadb-vector-database/>. วันที่ 25 มิถุนายน 2567

บทที่ 3

โมเดลสำหรับการประมวลผลข้อมูล

ในบทนี้จะกล่าวถึงการออกแบบโมเดล retrieval-augmented generation (RAG) ซึ่งเป็นกระบวนการที่ประกอบด้วย 2 ขั้นตอนหลัก ได้แก่ ขั้นตอนแรกเป็นการสร้าง “vector database” และขั้นตอนที่สองคือ “response generation”



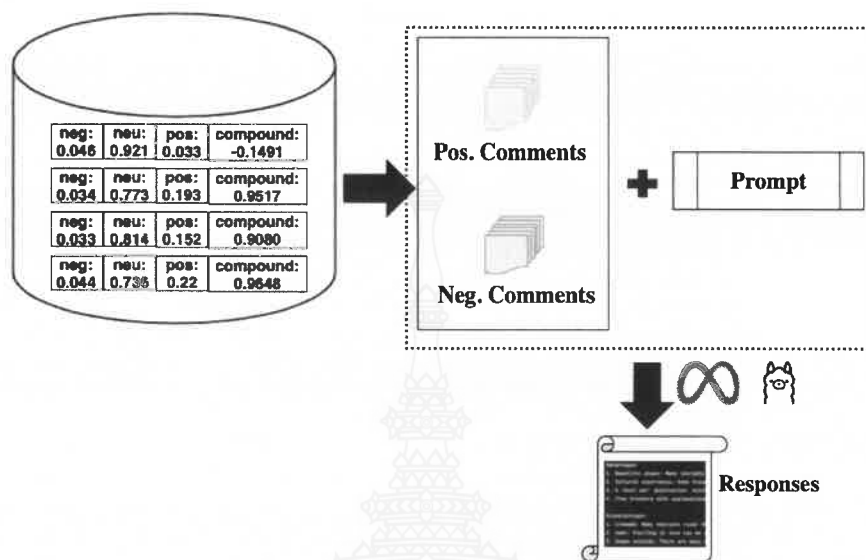
ภาพที่ 3.1: ขั้นตอนการสร้าง Vector Database

ตามที่แสดงในรูปที่ 3.1 “vector database” ทำหน้าที่เป็นแหล่งเก็บ “embedding” ซึ่งเป็นข้อมูลตัวเลข (numerical data) ที่อธิบายข้อมูลตัวหนังสือ (text data) ที่เก็บรวบรวมจากหน้าเว็บเพจ TripAdvisor

ข้อมูลที่เก็บจากหน้าเว็บเพจ TripAdvisor คือคำวิพากษ์วิจารณ์จากนักท่องเที่ยวที่เคยมาเที่ยวสถานที่ท่องเที่ยวสำคัญในกรุงเทพ โดยคำวิพากษ์วิจารณ์ที่เก็บรวบรวม สามารถแบ่งออกได้เป็น 3 กลุ่ม ได้แก่ กลุ่มที่ได้คะแนนจัดอันดับ (rating) 4-5 ดาว, อันดับ 3 ดาว, และอันดับ 1-2 ดาว

ขั้นตอนถัดไปเป็นการสร้างข้อมูลตัวเลข numerical data จากการวิเคราะห์อารมณ์ความรู้สึกของบทวิพากษ์วิจารณ์ (sentiment analysis) โดยคำวิพากษ์วิจารณ์แต่ละชิ้นจะถูกแปลงเป็นข้อมูลตัวเลข 4 ตัวแปร ได้แก่ neg., neu., pos. และ compound ซึ่งหมายถึง เปอร์เซ็นต์ที่บท

วิพากษ์วิจารณ์จะเป็นลบ เป็นกลาง เป็นบวก และบทสรุปในภาพรวม ตามลำดับ โดยบทวิพากษ์วิจารณ์หนึ่งจะเป็นบวกถ้าค่าของตัวแปร pos. มีค่ามากกว่าหรือเท่ากับ 0.05, เป็นลบเมื่อเท่ากับหรือน้อยกว่า -0.05



ภาพที่ 3.2: ขั้นตอนการสร้างคำตอบจาก RAG

ค่า compound กล่าวอีกอย่างหนึ่ง คือค่าตัวเลขที่เป็นตัวแทนของบทวิพากษ์วิจารณ์แต่ละชั้น และถูกจัดเก็บในฐานที่เป็น embedding ใน vector database และข้อมูล embedding เหล่านี้จะช่วยให้การสืบค้นข้อมูลออกจาก database ทำได้สะดวกและมีประสิทธิภาพ ดังแสดงในภาพที่ 3.2 ค่า compound value ทำหน้าที่เป็นดัชนี (index) สำหรับสืบค้นข้อมูล – ในกรณีต้องการสืบค้นข้อมูลบทวิพากษ์วิจารณ์เชิงบวก ตัวเลขดัชนีที่สืบค้นขึ้นมาจาก database จะต้องมีค่าตั้งแต่ 0.05 ขึ้นไป ในขณะที่การสืบค้นข้อมูลบทวิพากษ์วิจารณ์เชิงลบ ตัวเลขดัชนีต้องมีค่าเท่ากับ -0.05 หรือน้อยกว่าค่านี้

เอกสารที่ถูกสืบค้นขึ้นมาจาก database จะถูกรวมกันกับคำถามกระตุ้น (question prompt) ก่อนที่ข้อมูลทั้งหมดจะถูกส่งต่อให้ Llama2 – LLM ที่ประกอบอยู่ใน RAG เพื่อประมวลผลข้อมูลต่อไป โดยความคาดหวังแล้ว คำตอบ (response) จาก Llama2 จะแสดงให้เห็นในภาพสรุปเกี่ยวกับข้อดี (strength) และข้อจำกัด (weakness) ของการไปเที่ยวสถานที่ท่องเที่ยวสำคัญในกรุงเทพฯ ประเทศไทย

question = "Could you outline the advantages and disadvantages of visiting The Grand Palace in Bangkok, Thailand?"

prompt = ""

You are a bot that answers to a question from a user.

Your answer to this question MUST be based on comments of travellers who have visited The Grand Palace in Bangkok, Thailand.

Also, your answer MUST be in very short sentences and does not include extra information.

Here are the comments of travellers who have visited The Grand Palace in Bangkok, Thailand: {relevant_documents}

The question is: {question}

Compile a response to the user based on the travellers' comments and the question.

""

ภาพที่ 3.3: ตัวอย่างการสร้าง question prompt เพื่อป้อนให้กับ Llama2

ชุดของคำถามกระตุ้น (question prompt) และเอกสารบทวิพากษ์วิจารณ์ที่ถูกสร้างขึ้นก่อนที่จะป้อนให้กับ Llama2 นั้น ใช้หลักการของ “prompt engineering” ซึ่งในที่นี้ question prompt ประกอบด้วย 3 ส่วนหลัก ได้แก่

- 1) คำสั่งที่กำหนดหน้าที่สำหรับ Llama2 ประโยคที่แสดงส่วนนี้ (อ้างอิงภาพที่ 3.3) คือ “You are a bot that answers to a question from a user.”
- 2) คำถามที่กำหนดขอบเขตของคำตอบ (response) จาก Llama2 ประโยคที่แสดงส่วนนี้ (อ้างอิงภาพที่ 3.3) คือ “Your answer to this question MUST be based on comments of travelers who have visited The Grand Palace in Bangkok, Thailand” และส่วนของคำถาม “Could you outline the advantages and disadvantages of visiting The Grand Palace in Bangkok, Thailand?”
- 3) เอกสารแสดงบทวิพากษ์วิจารณ์ที่ถูกสืบค้นจาก vector database ประโยคที่แสดงส่วนนี้ (อ้างอิงภาพที่ 3.3) คือ “Here are the comments of travelers who have visited The Grand Palace in Bangkok, Thailand: {relevant_documents}”

เอาต์พุตจาก Llama2 แสดงถึงข้อดีที่สำคัญของสถานที่ท่องเที่ยว และข้อจำกัดที่ทำให้นักท่องเที่ยวรู้สึกไม่พอใจเมื่อเข้าเยี่ยมชมสถานที่นั้น ๆ ผลการประมวลผลข้อมูลทั้งหมดที่กล่าวมานี้ได้จากบทวิพากษ์วิจารณ์ทั้งในแง่บวก และแง่ลบที่เก็บรวบรวมไว้ใน vector database กระบวนการทดลอง และผลการทดลองซึ่งเป็นผลลัพธ์ที่ได้จากโมเดลที่อธิบายในบทนี้ จะปรากฏในบทที่ 4 และบทที่ 5 ตามลำดับ

บทที่ 4

การทดลอง และข้อมูลที่ใช้ในการทดลอง

ในบทนี้จะกล่าวถึงลักษณะและโครงสร้างของข้อมูลที่ใช้ในการทดลอง และกระบวนการทดลองในงานวิจัยฉบับนี้

ข้อมูลที่ใช้ในงานวิจัยนี้มาจากบทวิพากษ์วิจารณ์ของนักท่องเที่ยวบน TripAdvisor โดยการคัดเลือกบทวิพากษ์วิจารณ์มาจาก สถานที่ท่องเที่ยวสำคัญในกรุงเทพฯ ประเทศไทย ได้แก่ พระบรมมหาราชวัง (รวมถึงวัดพระศรีรัตนศาสดาราม) วัดเชตุพนวิมลมังคลาราม (หรือที่ประชาชนรู้จักโดยทั่วไปว่า วัดโพธิ์) วัดอรุณราชวรารามราชวรมหาวิหาร (หรือที่ประชาชนรู้จักโดยทั่วไปว่า วัดแจ้ง) ถนนข้าวสาร ตลาดนัดจตุจักร และเอเชียทีค เดอะ ริเวอร์ฟรอนท์ โดยในการค้นหาบทวิพากษ์วิจารณ์จะเป็นภาษาอังกฤษ เพื่อให้ได้บทวิพากษ์วิจารณ์เป็นภาษาอังกฤษที่โพสต์โดยนักท่องเที่ยวต่างชาติ โดยคำสำคัญที่ใช้ค้นหาสถานที่ท่องเที่ยวที่กล่าวมาข้างต้น ได้แก่ The Grand Palace, the Temple of the Reclining Buddha, the Temple of Dawn, Khaosan Road, Chatuchak Weekend Market และ Asiatique the Riverfront

ท่ามกลางข้อมูลมหาศาลบน TripAdvisor นักวิจัยพบว่าข้อมูลบางส่วนที่ล้าสมัยเกินไปที่จะนำมาสร้างบทวิเคราะห์ในงานวิจัยนี้ ในเงื่อนไขการเลือกข้อมูลจากหน้าเว็บจึงจำกัดช่วงเวลาการท่องเที่ยวในช่วงไตรมาสที่ 4 ปี 2566 หรือคือช่วงเดือนตุลาคม ถึงเดือนธันวาคมปี 2566 และในช่วงสองเดือนแรกของไตรมาสที่ 1 ปี 2567 ซึ่งก็คือช่วงเดือนมกราคม ถึงกุมภาพันธ์ปี 2567 ข้อดีของการเลือกช่วงเวลาการท่องเที่ยวในช่วง 5 เดือนที่กล่าวมาข้างต้น มีข้อดีคือ นอกจากจะเป็นข้อมูลการท่องเที่ยวล่าสุดแล้ว ยังเป็นข้อมูลที่อยู่ในช่วงวันหยุดคริสต์มาส และปีใหม่ที่มีนักท่องเที่ยวมาประเทศไทยมากที่สุดอีกด้วย

นอกเหนือจากที่นักท่องเที่ยวจะสามารถโพสต์คำวิพากษ์วิจารณ์บนหน้าเว็บของ TripAdvisor ได้แล้ว นักท่องเที่ยวยังสามารถให้คะแนน (rating) สถานที่และประสบการณ์จากการท่องเที่ยวได้ ดังแสดงในตารางที่ 4.1 ที่จำแนกจำนวนคำวิพากษ์วิจารณ์ตามระดับ rating 3 ระดับคือ ความพึงพอใจระดับสูง (High Preference) – คะแนน rating 4-5 ดาว, ความพึงพอใจระดับกลาง (Neutral Preference) – คะแนน rating 3 ดาว และความพึงพอใจระดับต่ำ (Low Preference) – คะแนน rating 1-2 ดาว

ตารางที่ 4.1: หมวดหมู่คำวิพากษ์วิจารณ์จาก TripAdvisor

Attractions	# Comments from TripAdvisor			Total
	High Preference (4 - 5 Stars)	Neutral Preference (3 Stars)	Low Preference (1 - 2 Stars)	
The Grand Palace	30 (60%)	11 (22%)	9 (18%)	50 (100%)
The Temple of the Reclining Buddha	20 (76.92%)	5 (19.23%)	1 (3.85%)	26 (100%)
The Temple of Dawn	25 (83.33%)	4 (13.33%)	1 (3.33%)	30 (100%)
Khaosan Road	17 (56.67%)	6 (20%)	7 (23.33%)	30 (100%)
Chatuchak Weekend Market	20 (68.97%)	6 (20.69%)	3 (10.34%)	29 (100%)
Asiatique The Riverfront	22 (61.11%)	10 (27.78%)	4 (11.11%)	36 (100%)
Total	134 (66.67%)	42 (20.90%)	25 (12.43%)	201 (100%)

ตามช่วงเวลา 5 เดือนที่กำหนด มีบทวิพากษ์วิจารณ์ทั้งสิ้น 201 ชิ้น จากสถานที่ท่องเที่ยวทั้งหมด 6 สถานที่ ที่มีรายละเอียดดังนี้

- 1) บทวิพากษ์วิจารณ์ 50 ชิ้น เกี่ยวกับ The Grand Palace – แบ่งเป็นบทวิพากษ์วิจารณ์ high preference, neutral preference, และ low preference จำนวน 30, 11 และ 9 ชิ้น ตามลำดับ
- 2) บทวิพากษ์วิจารณ์ 26 ชิ้น เกี่ยวกับ The Temple of the Reclining Buddha – แบ่งเป็นบทวิพากษ์วิจารณ์ high preference, neutral preference, และ low preference จำนวน 20, 5 และ 1 ชิ้น ตามลำดับ
- 3) บทวิพากษ์วิจารณ์ 30 ชิ้น เกี่ยวกับ The Temple of Dawn – แบ่งเป็นบทวิพากษ์วิจารณ์ high preference, neutral preference, และ low preference จำนวน 25, 4 และ 1 ชิ้น ตามลำดับ
- 4) บทวิพากษ์วิจารณ์ 30 ชิ้น เกี่ยวกับ Khaosan Road – แบ่งเป็นบทวิพากษ์วิจารณ์ high preference, neutral preference, และ low preference จำนวน 17, 6 และ 7 ชิ้น ตามลำดับ
- 5) บทวิพากษ์วิจารณ์ 29 ชิ้น เกี่ยวกับ Chatuchak Weekend Market – แบ่งเป็นบทวิพากษ์วิจารณ์ high preference, neutral preference, และ low preference จำนวน 20, 6 และ 3 ชิ้น ตามลำดับ

- 6) บทวิพากษ์วิจารณ์ 36 ชิ้น เกี่ยวกับ Asiatique The Riverfront – แบ่งเป็นบทวิพากษ์วิจารณ์ high preference, neutral preference, และ low preference จำนวน 22, 10 และ 4 ชิ้น ตามลำดับ

ข้อมูลวิพากษ์วิจารณ์จาก TripAdvisor ที่เก็บรวบรวมทั้งหมดตามจำนวนที่กล่าวข้างต้น ถูกนำมาวิเคราะห์อารมณ์ความรู้สึก (sentiment analysis) อีกครั้ง ด้วยเครื่องมือที่เรียกว่า VADER และได้ผลลัพธ์ ดังแสดงในตารางที่ 4.2

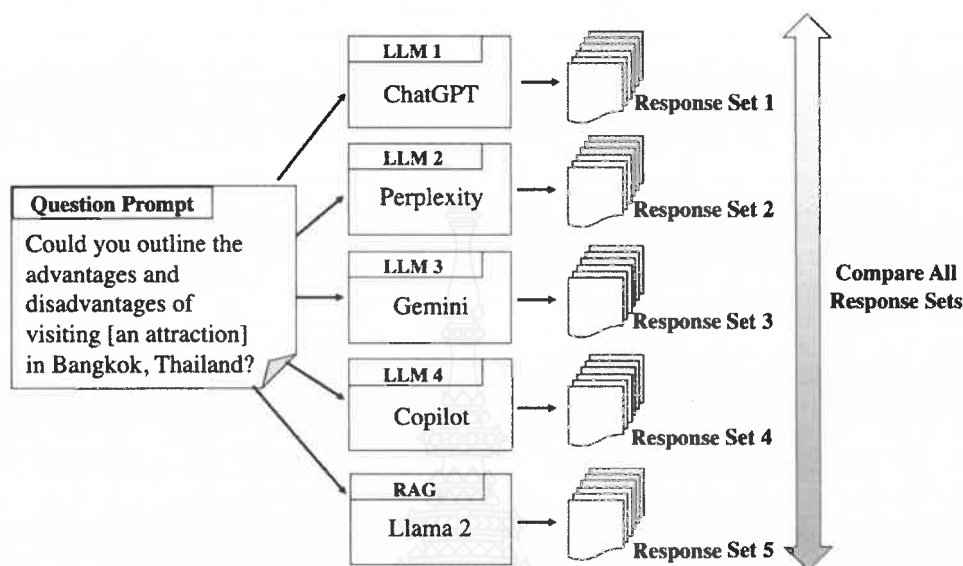
ตารางที่ 4.2: หมวดหมู่อารมณ์ความรู้สึกของบทวิพากษ์วิจารณ์จาก TripAdvisor

Attractions	Sentiment Categories			Total
	Positive	Neutral	Negative	
The Grand Palace	42 (84%)	0 (0%)	8 (16%)	50 (100%)
The Temple of the Reclining Buddha	25 (96.15%)	0 (0%)	1 (3.85%)	26 (100%)
The Temple of Dawn	28 (93.33%)	0 (0%)	2 (6.67%)	30 (100%)
Khaosan Road	25 (83.33%)	0 (0%)	5 (16.67%)	30 (100%)
Chatuchak Weekend Market	27 (93.10%)	0 (0%)	2 (6.90%)	29 (100%)
Asiatique The Riverfront	32 (88.89%)	1 (2.78%)	3 (8.33%)	36 (100%)
Total	179 (89.05%)	1 (0.50%)	21 (10.45%)	201 (100%)

ผลการวิเคราะห์ sentiment analysis แสดงให้เห็นว่าระดับ rating ตามที่ระบุในตาราง 4.1 มิได้แสดงการรับรู้ (perception) ของนักท่องเที่ยวได้ทั้งหมด ตัวอย่างเช่น ในตารางที่ 4.1 ระบุความพึงพอใจของนักท่องเที่ยวระดับ 4-5 ดาว จำนวน 66.67% แต่จากการวิเคราะห์ sentiment analysis พบว่า 89.05% ของนักท่องเที่ยวมีทัศนคติที่ดีต่อสถานที่ท่องเที่ยว

บทวิพากษ์วิจารณ์จากนักท่องเที่ยว ทั้งทัศนคติในแง่บวกและแง่ลบจะถูกใช้เป็นข้อมูลพื้นฐานสำหรับการวิเคราะห์ เพื่อหาบทสรุปเกี่ยวกับสถานที่ท่องเที่ยวทั้ง 6 สถานที่ในกรุงเทพฯ ดังที่กล่าวไปข้างต้น โดยที่ข้อมูลเหล่านี้จะใช้เป็นส่วนหนึ่งของกรอบงาน (framework) ที่ RAG ใช้ และจะได้ผลลัพธ์ที่คาดหวังคือ ข้อดีและข้อจำกัดของแต่ละสถานที่ ตามความคิดเห็นของนักท่องเที่ยว

การทดลองในงานวิจัยนี้ ต้องการแสดงให้เห็นสมรรถนะเปรียบเทียบระหว่าง RAG ซึ่งเป็นโมเดลที่นำเสนอในงานวิจัยนี้ และ LLMs ทั้ง 4 ตัวได้แก่ ChatGPT, Perplexity, Gemini และ Copilot



ภาพที่ 4.1: กระบวนการทดลองของงานวิจัยในภาพรวม

ภาพที่ 4.1 แสดงให้เห็นถึงภาพรวมของการทดลองสำหรับงานวิจัยนี้ ซึ่งสามารถอธิบายในเชิงองค์ประกอบของงานวิจัยได้ 3 องค์ประกอบ ได้แก่ ตัวแปรต้น (independent variable) ตัวแปรตาม (dependent variable) และตัวแปรควบคุม (control variable) โดยที่ตัวแปรต้นในที่นี้หมายถึง LLMs ทั้ง 4 ตัวคือ ChatGPT, Perplexity, Gemini และ Copilot และ RAG อีกหนึ่งตัวที่มีการออกแบบโดยใช้ training dataset ที่แตกต่างจาก LLMs ทั้ง 4 ตัว และภายใน RAG ใช้ Llama2 เป็น LLM สำหรับประมวลผล

สำหรับตัวแปรตาม ในงานวิจัยนี้คือ คำตอบ (response) ของ question prompt ซึ่งถูกพิจารณาเป็นตัวแปรควบคุม เนื่องจากทุก ๆ โมเดล – ทั้ง LLMs และ RAG จะถูกป้อน question prompt เดียวกันคือ “Could you outline the advantages and disadvantages of visiting [an attraction] in Bangkok, Thailand?”

เอาต์พุตจาก LLMs และ RAG ถูกเรียกว่า response ซึ่งแบ่งออกเป็นชุด sets 1 – 5 เนื่องจากการใช้โมเดลในการประมวลผลที่ต่างกัน response ที่ได้จึงแตกต่างกัน ในการทดลองจะทำการเปรียบเทียบ response ทั้ง 5 sets ตามที่ปรากฏตามภาพที่ 4.1 โดยผลลัพธ์จากการเปรียบเทียบ จะสามารถตอบคำถามวิจัย (research question) ต่อไปนี้

คำถามวิจัยที่ 1: อะไรคือข้อมูลใหม่ที่ถูกเปิดเผยโดย RAG และไม่เคยระบุโดย LLMs

การตอบคำถามวิจัยในข้อนี้ จะทำให้เห็นข้อจำกัดของ LLMs ที่ถูกแก้ไขได้โดยใช้ความสามารถของ RAG ซึ่งข้อจำกัดในแง่ของ LLMs เกิดจากข้อมูล training data ของ LLMs ขาดตกบกพร่องหรือมีความลำเอียง จึงส่งผลให้ผลลัพธ์ของการประมวลผลขาดประสิทธิภาพ

คำถามวิจัยที่ 2: ข้อมูลในส่วนใดที่แสดงถึงความสอดคล้องระหว่าง RAG และ LLMs

คำตอบของคำถามวิจัยข้อนี้ แสดงให้เห็นถึงความสม่ำเสมอของคำตอบของทั้ง RAG และ LLMs ซึ่งแสดงความน่าเชื่อถือของคำตอบจาก LLMs ในส่วนที่ไม่ใช่ข้อบกพร่องและไม่ลำเอียง

คำถามวิจัยที่ 3: ข้อมูลส่วนใดที่แสดงโดย LLM บางตัว แต่ไม่สอดคล้องกับ LLMs ตัวอื่น และไม่สอดคล้องกับ RAG ด้วยเช่นกัน

คำตอบของคำถามวิจัยในข้อนี้ ชี้ให้เห็นถึงแนวโน้มการเกิด hallucination ซึ่งทำให้ข้อมูลจาก LLM ตัวใดตัวหนึ่ง ไม่สอดคล้องกับ LLMs ตัวอื่น ๆ และไม่สอดคล้องกับ RAG

รายละเอียดของคำถามวิจัยข้อที่ 1 – 3 จะแสดงในส่วนของผลการทดลองในบทที่ 5 ต่อไป



บทที่ 5

ผลการทดลอง

ข้อมูลในบทนี้จะนำเสนอผลลัพธ์ response จาก RAG และ LLMs ทั้ง 4 ตัวได้แก่ ChatGPT, Perplexity, Gemini และ Copilot ข้อเปรียบเทียบ response ของโมเดลทั้งหมดนี้ ทำให้เห็นถึงความน่าเชื่อถือของโมเดล ดังแสดงในตารางที่ 5.1

ผลลัพธ์ response จากทั้ง 5 โมเดล – รวมถึงทั้ง LLMs และ RAG คือคำตอบของ question prompt “Could you outline the advantages and disadvantages of visiting the Grand Palace in Bangkok, Thailand?” ที่ได้รับการจัดหมวดหมู่ (แสดงในตารางที่ 5.1) ว่ามีทั้งหมด 20 responses ทั้งนี้ response จาก RAG จะถูกแบ่งออกเป็น 2 ส่วนคือ RAG-POS และ RAG-NEG ซึ่งหมายถึง response จาก RAG ที่ประมวลผลบนข้อมูลวิพากษ์วิจารณ์ของนักท่องเที่ยวที่เป็น positive และ negative ตามลำดับ

เอาต์พุตทั้งหมดจาก LLMs และ RAG คือ response ลำดับที่ 1 – 20 (ดังแสดงในตารางที่ 5.1) response ลำดับที่ 1 – 7 แสดงบทสรุปของข้อดี (advantages) ของการเข้าเยี่ยมชม The Grand Palace ตามความเห็นของนักท่องเที่ยว และลำดับที่ 8 – 20 แสดงบทสรุปข้อจำกัด (disadvantages) ของสถานที่ท่องเที่ยวแห่งนี้ ตามลำดับ จาก response ทั้งหมด 20 หัวข้อ สามารถสรุปได้ ดังต่อไปนี้

- 1) ความสม่ำเสมอของ response จาก RAG และ LLMs แสดงให้เห็นถึงความน่าเชื่อถือของ response ทั้งจาก RAG และ LLMs จากตัวอย่างในตารางที่ 5.1 response ในส่วนนี้คือ response หมายเลข 1, 2, 5, 8, 11 และ 16 กล่าวโดยสรุป response ในส่วนนี้แสดงให้เห็นว่าถึงแม้จะใช้โมเดลที่แตกต่างกันในการประมวลผลข้อมูล จะมีบทสรุปจากการประมวลผลส่วนหนึ่งที่ทั้ง RAG และ LLMs เห็นตรงกัน
- 2) จำเพาะบาง response จาก RAG ที่ไม่เคยเปิดเผยโดย LLMs แสดงให้เห็นความน่าเชื่อถือของ RAG ที่มากกว่า LLMs เนื่องจาก response กลุ่มนี้ถูกประมวลผลมาจากข้อมูลใน vector database ซึ่งมีความใหม่มากกว่าข้อมูล training data ของ LLMs หรือกล่าวอีกนัยหนึ่งคือ response ของ RAG ได้รับการรับรองความน่าเชื่อถือโดยข้อมูลใน vector database ซึ่งข้อมูลในส่วนนี้แสดงให้เห็นถึงส่วนบกพร่องของ response ที่ได้จาก LLMs จากตัวอย่างในตารางที่ 5.1 response ในส่วนนี้คือ response หมายเลข 4, 6, 12, 13, 14, 15 และ 17

ตารางที่ 5.1: เอาต์พุต Response จาก RAG และ LLMs

Response Aspect No.	Response Aspects	RAG-POS	RAG-NEG	ChatGPT	Perplexity	Gemini	Copilot
1	Beautiful architecture and design	X	X	X	X	X	X
2	Rich historical and cultural significance	X		X	X	X	X
3	Sacred site, called Temple of the Emerald Buddha, housed within the Grand Palace			X	X	X	X
4	Considered a must-see destination	X					
5	Offers good photo opportunities	X			X		
6	Free brochures with explanations available		X				
7	Situated in the heart of Bangkok			X		X	
8	Crowded and touristy	X	X	X	X	X	X
9	Weather that is hot and humid					X	X
10	Strict dress code			X	X	X	X
11	Expensive entry fee	X		X			
12	Limited information provided	X					
13	Encounter with disrespectful staff	X	X				
14	Perceived unsafe environment for belongings	X					
15	Lack of signs or posters for guidance		X				
16	Presence of scammers outside		X	X	X		X
17	Perception of being overpriced and under-maintained		X				
18	Limited access during royal ceremonies or events			X		X	
19	Lack of novelty				X		
20	Mobility challenges						X

- 3) ความสม่ำเสมอของ response ของ LLMs ตั้งแต่ 2 ตัวขึ้นไป ซึ่งไม่ปรากฏใน response ของ RAG ถือว่าเป็นข้อบกพร่อง (error) ของการทำงานของ RAG เหตุการณ์ลักษณะนี้เกิดขึ้นเมื่อ response จาก LLMs มีความสอดคล้องกัน และสอดคล้องกับข้อมูลใน vector database แต่ผลเอาต์พุต response จาก RAG ไม่ปรากฏข้อมูลในส่วนนั้น จึงถือเป็น error ของข้อมูลที่ผลิตโดย RAG ตัวอย่างที่ปรากฏในตารางที่ 5.1 คือ response หมายเลข 3, 7, 9, 10 และ 18
- 4) จำเพาะ response จาก LLM ตัวใดตัวหนึ่งที่ไม่ปรากฏใน response ของ LLMs ตัวอื่น และ response จาก RAG หรือเป็น response ที่ขัดแย้งไม่สอดคล้องกับ LLMs และ RAG แล้วนั้น ถือว่าเป็น error ของ LLM ตัวนั้น ตัวอย่างในตารางที่ 5.1 คือ response หมายเลข 19 และ 20 ซึ่ง error เหล่านี้ ถูกค้นพบว่าเป็นผลจากการประมวลผลข้อมูลที่ผิดพลาด หรือใช้ข้อมูลที่ไม่เป็นปัจจุบัน ตัวอย่างเช่น Copilot แสดง response ลำดับที่ 20 (แสดงในตารางที่ 5.1) ที่ว่า พื้นที่โดยรอบพระบรมมหาราชวัง มีความยากลำบากสำหรับผู้พิการ และคนชราที่ต้องใช้รถเข็น (wheel chair) ในความจริงแล้ว เป็นความผิดพลาดของการประมวลผลข้อมูลของ Copilot เนื่องจากข้อมูลที่ Copilot ใช้อ้างอิงแท้จริงได้กล่าวไว้ว่า ณ ด้านหน้าบริเวณจุดขายบัตรมีรถเข็นให้บริการสำหรับผู้พิการ และคนชรา แต่การตีความของ Copilot เป็นเพียงการคาดเดาปัญหาที่อาจเกิดขึ้น ในที่นี้จึงทำให้เกิด error ประเภท hallucination

ในภาพรวม จากบทวิเคราะห์ในข้อ 1) ถึง 4) แสดงให้เห็นว่าผลการทำงานของ LLMs มีทั้งจุดแข็ง ที่แสดงให้เห็นถึงระดับความน่าเชื่อถือของโมเดล และข้อจำกัดที่ทำให้ความน่าเชื่อถือของโมเดลลดลง เพื่อให้การเปรียบเทียบระหว่างโมเดลเข้าใจได้ง่ายขึ้น นักวิจัยจึงขอเสนอมาตรวัดความน่าเชื่อถือ “reliability metric” ดังแสดงในสมการที่ 1 (Eq. 1)

$$\mathcal{R} = \frac{n}{u} \quad (\text{Eq. 1})$$

สมการ (Eq. 1) อธิบายความน่าเชื่อถือ ($\mathcal{R}$) ว่าเป็นอัตราส่วนระหว่าง จำนวน response ที่น่าเชื่อถือจาก RAG (n_{RAG}) หรือจำนวน response ที่น่าเชื่อถือจาก LLM (n_{LLM}) ต่อจำนวน response ที่น่าเชื่อถือทั้งหมด (u) – ทั้งที่เป็น response จากทั้ง RAG และ LLM

จำนวน response ที่น่าเชื่อถือจาก RAG (n_{RAG}) หมายถึง จำนวน response จาก RAG ที่สะท้อนข้อมูลใน vector database ได้อย่างถูกต้อง จากตัวอย่างข้อมูลในตารางที่ 5.1

n_{RAG} เท่ากับ 13 (แบ่งออกเป็นจำนวน response จาก $n_{RAG-POS} = 9$ และ $n_{RAG-NEG} = 4$) ในจำนวนนี้คือ response หมายเลข 1, 2, 4, 5, 6, 8, 11, 12, 13, 14, 15, 16 และ 17

จำนวน response ที่น่าเชื่อถือจาก LLM (n_{LLM}) หมายถึง จำนวน response จาก LLM ตัวหนึ่งที่สอดคล้องกับ response ของ LLMs อื่น หรือสอดคล้องกับ response จาก RAG จากตัวอย่างข้อมูลในตารางที่ 5.1 สามารถสรุปได้ว่า n_{LLM} ของ ChatGPT, Perplexity, Gemini และ Copilot คือ 9, 7, 8 และ 7 ตามลำดับ

จำนวน response ที่น่าเชื่อถือทั้งหมดจากทั้ง RAG และ LLMs (U) คือจำนวน response ที่ถูกต้องและน่าเชื่อถือทั้งหมด จากตารางที่ 5.1 ค่า U เท่ากับ 18 ซึ่งหมายถึง response ลำดับที่ 1 ถึง 18 (response ลำดับที่ 19 และ 20 คือ hallucination ของ Perplexity และ Copilot)

ตารางที่ 5.2: ค่า Reliability ของ RAG และ LLMs

Models	Reliability ($\mathcal{R}$)
RAG	0.72
ChatGPT	0.50
Perplexity	0.39
Gemini	0.44
Copilot	0.39

ตารางที่ 5.2 แสดงค่าความน่าเชื่อถือ ($\mathcal{R}$) ของ RAG และ LLMs โดยคำนวณจากข้อมูลที่ปรากฏในตารางที่ 5.1 ซึ่งเป็น response ที่ตอบคำถามเกี่ยวกับ The Grand Palace โดยค่าความน่าเชื่อถืออยู่ระหว่าง $[0, 1]$ โดยค่า 0 ระบุความน่าเชื่อถือที่ต่ำที่สุด และ 1 ระบุความน่าเชื่อถือที่สูงที่สุด สำหรับ response ของข้อมูลในตารางที่ 5.1 ค่าความน่าเชื่อถือของ RAG, ChatGPT, Gemini, Perplexity และ Copilot เท่ากับ 0.72, 0.50, 0.44 และ 0.39 (มีค่าเท่ากันสำหรับ Perplexity และ Copilot) ตามลำดับ

ตารางที่ 5.2 แสดงค่าความน่าเชื่อถือ ($\mathcal{R}$) ของ response จาก RAG และ LLMs ที่อธิบายเกี่ยวกับ The Grand Palace โดยค่าความน่าเชื่อถือนี้อยู่ถูกคำนวณกับ response จาก RAG และ LLMs ที่อธิบายสถานที่ท่องเที่ยวอื่น ๆ ได้แก่ The Temple of the Reclining Buddha, The

Temple of Dawn, Khaosan Road, Chatuchak Weekend Market และ Asiatique the Riverfront ตารางที่ 5.3 แสดงบทสรุปของค่าความน่าเชื่อถือ n_{LLM} , n_{RAG} , U และ $\mathcal{R}$ ของทุกโมเดล

ตารางที่ 5.3: ค่า Reliability จำแนกตามบทวิพากษ์วิจารณ์ของแต่ละสถานที่ท่องเที่ยว

Attractions	# All Reliable Responses	#Reliability Metrics	Models				
			RAG	ChatGPT	Perplexity	Gemini	Copilot
The Grand Palace	$U = 18$	n	13	9	7	8	7
		$\mathcal{R}$	0.72	0.5	0.39	0.44	0.39
The Temple of the Reclining Buddha	$U = 14$	n	11	6	5	8	4
		$\mathcal{R}$	0.79	0.43	0.36	0.57	0.29
The Temple of Dawn	$U = 16$	n	12	7	7	8	6
		$\mathcal{R}$	0.75	0.44	0.44	0.50	0.38
Khaosan Road	$U = 18$	n	11	9	9	8	8
		$\mathcal{R}$	0.61	0.5	0.5	0.44	0.44
Chatuchak Weekend Market	$U = 18$	n	11	9	8	7	5
		$\mathcal{R}$	0.61	0.5	0.44	0.39	0.28
Asiatique the Riverfront	$U = 13$	n	12	9	4	7	5
		$\mathcal{R}$	0.92	0.69	0.31	0.54	0.38

ข้อมูลในตารางที่ 5.3 แสดงค่าความน่าเชื่อถือของ response เกี่ยวกับ The Temple of the Reclining Buddha จาก RAG และ LLMs – ChatGPT, Perplexity, Gemini และ Copilot ข้อมูลในตารางระบุว่า RAG มีประสิทธิภาพกว่า LLMs ทุกตัว โดยค่าความน่าเชื่อถือของ RAG เท่ากับ 0.79 ซึ่งสูงกว่าค่าความน่าเชื่อถือของ LLMs ที่มีค่าระหว่าง 0.29 ถึง 0.57

ค่าความน่าเชื่อถือของ response เกี่ยวกับ The Temple of Dawn จาก RAG เท่ากับ 0.75 ซึ่งสูงกว่าค่าความน่าเชื่อถือของ response จาก LLMs ที่มีค่าระหว่าง 0.38 ถึง 0.50 ค่าความน่าเชื่อถือของ response เกี่ยวกับ Khaosan Road จาก RAG เท่ากับ 0.61 ซึ่งสูงกว่าค่าความน่าเชื่อถือของ response จาก LLMs ที่มีค่าระหว่าง 0.44 ถึง 0.50 ค่าความน่าเชื่อถือของ response เกี่ยวกับ Chatuchak Weekend Market จาก RAG เท่ากับ 0.61 ซึ่งสูงกว่าค่าความน่าเชื่อถือของ response จาก LLMs ที่มีค่าระหว่าง 0.28 ถึง 0.50 สุดท้ายค่าความน่าเชื่อถือของ response เกี่ยวกับ Asiatique the Riverfront จาก RAG เท่ากับ 0.92 ซึ่งสูงกว่าค่าความน่าเชื่อถือของ response จาก LLMs ที่มีค่าระหว่าง 0.31 ถึง 0.69

เป็นที่น่าสนใจว่าความน่าเชื่อถือของ RAG สำหรับ Khaosan Road และ Chatuchak Weekend Market ลดลงเล็กน้อย เมื่อมีการเปรียบเทียบกับข้อมูลชุดที่วิเคราะห์จากสถานที่อื่น เนื่องจากข้อมูลวิพากษ์วิจารณ์ของสถานที่ท่องเที่ยวทั้งสองที่ค่อนข้างกว้าง ทำให้เกิดความยากลำบากในการสรุปข้อมูลที่ถูกต้องแม่นยำโดย RAG แต่แม้กระนั้น ประสิทธิภาพการประมวลผลข้อมูลของ RAG ยังดีกว่าที่ LLMs สามารถทำได้



บทที่ 6

บทวิเคราะห์และงานวิจัยสืบเนื่องในอนาคต

บทนี้จะกล่าวถึงบทวิเคราะห์จากผลลัพธ์จากงานวิจัย โดยบทวิเคราะห์จะแบ่งเป็นประเด็นตามคำถามวิจัยทั้ง 3 ข้อ ดังที่ได้กล่าวไว้ในบทที่ 4 ได้แก่ (1) ข้อมูลที่ค้นพบใหม่โดย RAG และไม่เคยกล่าวถึงในเอาต์พุตจาก LLMs (2) ข้อมูลที่สอดคล้องกันระหว่างเอาต์พุตจาก RAG และ LLMs (3) ข้อมูลเอาต์พุตจาก LLMs ที่แสดงถึงข้อผิดพลาดและการเกิด hallucination ของ LLMs

ในประเด็นของงานวิจัยสืบเนื่องในอนาคตจะเป็นคำแนะนำที่รัฐบาล และองค์การบริหารส่วนท้องถิ่นสามารถนำไปใช้เพื่อสร้างเทศบัญญัติ (local ordinance) และข้อบังคับ (bylaw) ที่ใช้เป็นแนวปฏิบัติสำหรับบุคคลที่เกี่ยวข้องเพื่อให้สถานที่ท่องเที่ยวที่น่าสนใจสำหรับนักท่องเที่ยวมากยิ่งขึ้น

คำตอบของคำถามวิจัยข้อที่ 1: ข้อมูลที่ถูกค้นพบโดย RAG สรุปได้ถึงสิ่งที่นักท่องเที่ยวมองหาจากการท่องเที่ยวในประเทศไทย ดังนี้

- 1) **เอกลักษณ์ของสถานที่ท่องเที่ยว:** นักท่องเที่ยวคาดหวังที่จะเห็นเอกลักษณ์เฉพาะของประเทศไทยและสถานที่ที่เข้าเยี่ยมชม ตัวอย่างเช่น นักท่องเที่ยวที่เข้าชมพระบรมมหาราชวัง (The Grand Palace) และวัดต่าง ๆ ในประเทศไทย ต้องการศึกษาศิลปวัฒนธรรม และศิลปะที่เกี่ยวข้องกับศาสนาและความเชื่อของคนในประเทศ ในขณะเดียวกันการท่องเที่ยวในตลาดพื้นเมืองอย่างตลาดนัดจตุจักร (Chatuchak Weekend Market) และเอเชียทีค เดอะ ริเวอร์ฟรอนท์ (Asiatique The Riverfront) นักท่องเที่ยวคาดหวังเป็นอย่างมากที่จะได้รับประสบการณ์เกี่ยวกับอาหารท้องถิ่นแท้ สตรีทฟู้ด (street foods) และนวดแผนไทย
- 2) **ข้อมูลเกี่ยวกับสถานที่ท่องเที่ยว:** นักท่องเที่ยวต้องการเที่ยวแบบมีข้อมูล และรับรู้ความหมายของสิ่งที่เยี่ยมชม ไม่ว่าจะเป็นตึก รูปปั้น อนุสาวรีย์ ภาพวาด ประวัติศาสตร์ งานศิลป์ วัฒนธรรม และค่านิยม ทุกอย่างที่พบเห็นจะมีคุณค่าหากนักท่องเที่ยวเข้าใจสิ่งที่พบเห็น ซึ่งข้อมูลอาจมาในรูปแบบสื่อสิ่งพิมพ์ ไกด์เสียง (audio tour guide) แม้กระทั่งป้าย สัญลักษณ์ โปสเตอร์ ที่ช่วยให้นักท่องเที่ยวสามารถท่องเที่ยวและเข้าใจความหมายของสิ่งต่าง ๆ ได้ด้วยตนเอง
- 3) **การบริการและความสะดวกสบาย:** คุณภาพของการบริการและความสะดวกสบายมีความสำคัญต่อการสร้างประสบการณ์ให้นักท่องเที่ยว พนักงานที่มีทัศนคติให้ความเคารพและเต็มใจให้บริการสามารถสร้างความพึงพอใจและความประทับใจให้นักท่องเที่ยวได้มาก ตัวอย่างจากคำ

วิพากษ์วิจารณ์พบว่า สิ่งเล็กน้อยเช่น การให้น้ำดื่มฟรีกับนักท่องเที่ยว หรือการบริหารจัดการให้นักท่องเที่ยวได้เข้าชมสถานที่ ในวันที่อากาศร้อนมีนักท่องเที่ยวหนาแน่น ทำให้นักท่องเที่ยวเกิดความประทับใจโดยเฉพาะอย่างยิ่งกลุ่มนักท่องเที่ยวที่มีความกังวลต่อการแพร่เชื้อโรค

- 4) **วิธีการชำระเงิน:** นักท่องเที่ยวต้องการให้มีวิธีการจ่ายเงินที่หลากหลาย โดยเฉพาะเมื่อมีการซื้อบัตรเข้าชมสถานที่ต่าง ๆ ซึ่ง ณ ปัจจุบันสามารถชำระได้ด้วยวิธีการใช้เงินสดเพื่อซื้อบัตรหน้าประตูทางเข้าสถานที่เท่านั้น
- 5) **ความปลอดภัย:** นักท่องเที่ยวบางกลุ่มพบปัญหาการหลอกลวงโดยไกด์นำเที่ยวที่ไม่ถูกกฎหมาย และพบปัญหาโดนล้วงกระเป๋า จึงมีคำขอจากนักท่องเที่ยวให้มีการตรวจท่องเที่ยวประจำสถานที่ โดยเฉพาะช่วงเวลาที่มียกนักท่องเที่ยวแออัดคับคั่ง
- 6) **ข้อมูลเวลาเปิดทำการ:** นักท่องเที่ยวจำนวนมากขอให้มีการให้ข้อมูลที่ชัดเจนเกี่ยวกับช่วงเวลาเปิด-ปิดทำการของสถานที่ต่าง ๆ และช่วงเวลาที่เหมาะสมที่สุดในการเข้าเยี่ยมชม เนื่องจากบางสถานที่ เช่น Khaosan Road และ Asiatique the Riverfront มีร้านค้าเปิดให้บริการทั้งกลางวันและกลางคืน แต่ช่วงเวลาที่การแสดง และเหตุการณ์สำคัญเกิดขึ้นมักเกิดในช่วงเย็นถึงช่วงค่ำ
- 7) **ความหลากหลายของสินค้า ของที่ระลึก และอาหาร:** นักท่องเที่ยวมีความคาดหวังว่าจะเห็นและได้เลือกซื้อสินค้า ของที่ระลึก และอาหารที่เป็นเอกลักษณ์ของไทยแท้ และยังคงคาดหวังที่จะเห็นความหลากหลายของสินค้าเมื่อเยี่ยมชมสถานที่ที่แตกต่างกัน นอกจากนี้ยังคาดหวังว่า ในสถานที่ท่องเที่ยวเดียวกัน ร้านค้าที่ขายของควรมีความหลากหลายของสินค้าและอาหารให้ได้เลือกซื้อ
- 8) **ความเป็นมิตรของคนท้องถิ่น:**เสน่ห์ของคนไทยที่สำคัญคือ ความเป็นมิตรไมตรีทั้งจากผู้ค้าและประชาชนพื้นเมือง จึงนับเป็นหนึ่งในจุดแข็งที่ทำให้ให้นักท่องเที่ยวรู้สึกประทับใจ และมีประสบการณ์ที่ดีในการท่องเที่ยวมากที่สุด
- 9) **ข้อมูลประกอบการตัดสินใจในการเลือกสถานที่ท่องเที่ยว:** เนื่องจากสถานที่ท่องเที่ยวบางแห่งมีข้อมูลบางอย่างที่อาจไม่ได้ระบุในเอกสารโปรโมทการท่องเที่ยวอย่างเป็นทางการ แต่มีความจำเป็นที่จะต้องแจ้งให้นักท่องเที่ยวทราบล่วงหน้าเพื่อประกอบการตัดสินใจก่อนมาเที่ยว เช่น Khaosan Road ที่เป็นสถานที่ที่เหมาะสมกับการเที่ยวกลางคืน แต่อาจไม่เหมาะสมกับนักท่องเที่ยวที่มีครอบครัว หรือมีสมาชิกในครอบครัวที่เป็นเด็กเล็กร่วมเดินทางมาด้วย หรือกรณีการเข้าเยี่ยมชม The Grand Palace นักท่องเที่ยวควรได้รับข้อมูลเพิ่มเติมหากบางสถานที่ภายในปิดเพื่อปรับปรุงหรือปิดเนื่องจากมีพระราชพิธีสำคัญที่ไม่อนุญาตให้นักท่องเที่ยวมีส่วนร่วม

- 10) **ความเป็นส่วนตัว:** แม้แต่ในสถานที่สาธารณะ นักท่องเที่ยวก็ต้องการความเป็นส่วนตัว การอนุญาตให้มีสตูดิโอถ่ายภาพนอกสถานที่ เป็นตัวอย่างหนึ่งที่นักท่องเที่ยวรู้สึกถูกรบกวนจากการเข้าเยี่ยมชม ที่ต้องการความเงียบสงบเพื่อซึมซับความงามของสถานที่นั้น

คำตอบของคำถามวิจัยข้อที่ 2: ข้อมูลได้จาก RAG และ LLMs มีความสอดคล้องกันบางประการ ซึ่งข้อมูลในส่วนนี้ถือว่าน่าเชื่อถือ และเป็นประโยชน์ที่จะนำไปพัฒนาต่อยอดเพื่อสร้างคุณค่าให้กับสถานที่ท่องเที่ยว ดังนี้

- 1) **ความหนาแน่นของนักท่องเที่ยว** ข้อมูลจาก RAG และ LLMs สอดคล้องกันในประเด็นที่ว่า สถานที่ท่องเที่ยวที่สำคัญในกรุงเทพฯ มีความหนาแน่นมาก โดยเฉพาะในช่วงฤดูท่องเที่ยว (high season) อย่างวันหยุดยาวช่วงคริสต์มาสและปีใหม่
- 2) **อากาศร้อน** ภาวะอากาศร้อนอบอ้าว เป็นสิ่งที่หลีกเลี่ยงไม่ได้เมื่อท่องเที่ยวในประเทศไทย ทั้ง RAG และ LLMs ให้ข้อมูลที่ตรงกันในแง่ที่ว่า สถานที่ท่องเที่ยวในประเทศอยู่กลางแจ้ง ทำให้หลีกเลี่ยงการปะทะกับอากาศร้อน และแสงแดดได้ จากข้อมูลที่ให้โดย RAG สามารถสรุปได้ว่า นักท่องเที่ยวที่เคยเที่ยวในประเทศไทยมักแนะนำว่า ช่วงเวลาที่แนะนำสำหรับสถานที่ท่องเที่ยว ยอดนิยมอย่าง The Grand Palace คือช่วงก่อนเที่ยงที่อากาศยังไม่ร้อนมาก
- 3) **ข้อบังคับการแต่งกาย** สำหรับสถานที่ศักดิ์สิทธิ์ ได้แก่ The Grand Palace (สถานที่ตั้ง The Emerald Buddha Temple), The Temple of the Reclining Buddha และ The Temple of Dawn นักท่องเที่ยวจะต้องแต่งกายคลุมไหล่และหัวเข่า ข้อมูลจากทั้ง RAG และ LLMs ระบุตรงกันว่าข้อบังคับเกี่ยวกับการแต่งกายควรมีความยืดหยุ่นได้บ้าง โดยอนุญาตให้นักท่องเที่ยวได้เลือกรูปแบบการแต่งกายที่สุภาพในแบบเฉพาะตัว และให้เหมาะสมกับสภาวะอากาศในประเทศไทยที่ร้อนจัด
- 4) **ค่าธรรมเนียมเข้าเยี่ยมชมสถานที่** ค่าธรรมเนียมค่าเข้าสถานที่ที่แพง ในที่นี้นักท่องเที่ยวมิได้พิจารณาเป็นค่าธรรมเนียมในแต่ละที่ แต่หมายถึงค่าธรรมเนียมสำหรับเข้าทุกสถานที่รวมกันมีราคาแพง กล่าวคือ การเยี่ยมชม “Rattanakosin Island” (เกาะรัตนโกสินทร์) หมายถึงการเยี่ยมชมสถานที่ท่องเที่ยวที่สำคัญของกรุงเทพฯ หลายที่ซึ่งนิยมทำภายในวันเดียวกัน ได้แก่ การเยี่ยมชม The Grand Palace, The Temple of the Reclining Buddha และ The National Museum Bangkok รวมถึงสถานที่ที่เดินทางไปถึงได้จากเกาะรัตนโกสินทร์โดยการนั่งเรือข้ามฟากอย่าง The Temple of Dawn ข้อมูลจาก RAG และ LLMs ระบุว่าการเยี่ยมชมสถานที่เหล่านี้สามารถทำได้ภายในวันเดียว แต่นักท่องเที่ยวรู้สึกว่าการเข้าชมสถานที่เหล่านี้รวมกันแล้ว

ค่อนข้างสูง นักท่องเที่ยวจึงมีความเห็นว่าจะเป็นการประหยัดงบประมาณได้มาก หากมีการเสนอขายบัตรเข้าชมรวมทุกสถานที่เข้าด้วยกัน

- 5) การเดินทาง ข้อมูลจาก RAG และ LLMs สรุปรวมกันว่า การเดินทางและระบบขนส่ง ส่งผลอย่างมากต่อความประทับใจที่นักท่องเที่ยวจะมีให้กับสถานที่ท่องเที่ยวหนึ่ง ตัวอย่างเช่น การเดินทางโดยเรือแท็กซี่ ที่รับ-ส่งฟรีให้กับ Asiatique The Riverfront สร้างความประทับใจให้นักท่องเที่ยวได้ แต่ในทางกลับกันการขาดแคลนรถสาธารณะอย่างรถแท็กซี่ในช่วงเวลาที่มีผู้ต้องการสูงสุด เช่น ช่วงเวลาหลังอาหารค่ำที่ Asiatique The Riverfront ก็ทำให้นักท่องเที่ยวรู้สึกหงุดหงิดใจได้
- 6) ราคาที่พัก สินค้า และอาหารที่แพงเกินจริงสำหรับนักท่องเที่ยว บทสรุปจาก LLMs และ RAG ตรงกันในประเด็นที่ว่า ณ สถานที่ท่องเที่ยวบางแห่ง เช่น Khaosan Road และ Asiatique the Riverfront ราคาที่พัก อาหาร และสินค้าที่ขายให้นักท่องเที่ยวจะแพงเกินความเป็นจริงเมื่อเทียบกับราคาขายในสถานที่อื่น
- 7) การก่อวินาศกรรมและการล่วงละเมิด พฤติกรรมที่ก้าวร้าวที่พบจากผู้ค้าตามถนน (street vendor) และผู้ให้บริการขนส่งสาธารณะ เช่น รถสามล้อ (รถตุ๊กตุ๊ก) และรถแท็กซี่ ที่พยายามกดดันเพื่อขายสินค้าและบริการให้นักท่องเที่ยว จากข้อมูลจาก RAG และ LLMs พบตรงกันว่ามักพบบริเวณ Khaosan Road

คำตอบของคำถามวิจัยข้อที่ 3: ข้อมูลได้จาก LLMs ได้ถูกนำมาวิเคราะห์ และพบว่า มีข้อบกพร่องและ hallucination อยู่ในบางคำตอบของ LLMs บางตัว จำแนกได้ดังนี้

- 1) ข้อบกพร่องเกิดจากการมีข้อมูลไม่ครบถ้วน เนื่องจาก LLM จะสร้าง response จาก training data อย่างไรก็ตาม หาก training data ที่ใช้สำหรับสอนโมเดลของ LLM มีไม่ครบถ้วน อาจทำให้เกิดกรณี response จาก LLM ที่คลาดเคลื่อนจากความเป็นจริง ตัวอย่างเช่น ในการทดลองของงานวิจัยนี้พบว่า response จาก ChatGPT ระบุว่าสถานที่ของ Asiatique the Riverfront เป็นสถานที่เปิด ซึ่งมีจุดอ่อน เนื่องจากความสนุกสนานของนักท่องเที่ยวขึ้นอยู่กับสภาพอากาศ – หากฝนตกหนักหรืออากาศร้อนจัด กิจกรรมต่าง ๆ ทั้งหมดอาจต้องยกเลิก

การสร้าง response จาก ChatGPT ตามที่กล่าวข้างต้นนั้นถูกต้อง ในแง่ที่นักท่องเที่ยวต้องการดูวิวทิวทัศน์ของแม่น้ำเจ้าพระยาและต้องการจองโต๊ะสำหรับรับประทานอาหารด้านนอกร้านอาหาร อย่างไรก็ตามในกรณีที่อากาศไม่ดี นักท่องเที่ยวยังคงสามารถรับประทานอาหารภายในร้าน ดูโชว์ และช้อปปิ้งภายในอาคาร

- 2) **ความขัดแย้งของข้อมูล** เป็นสาเหตุให้เกิด hallucination ของ LLM ซึ่งแบ่งออกเป็น 2 ลักษณะ กรณีแรก hallucination ถูกสังเกตได้จาก ข้อมูลใน response เดียวกันมีความขัดแย้งกัน และ ในกรณีที่สอง สังเกตได้จาก ข้อความจาก response ที่คนละชั้นที่ให้ข้อมูลที่แตกต่างตรงกันข้าม ตัวอย่างจากการทดลองพบว่า Perplexity รายงานใน response เดียวกันว่า ข้อดีของ Asiatique the Riverfront คือมีร้านอาหารที่หลากหลาย ในขณะที่ response เดียวกันรายงานว่า ข้อเสียคือ มีอาหารให้เลือกจำกัด ซึ่งข้อมูลที่ให้ทั้งสองชุดใน response เดียวกันค่อนข้างมีความขัดแย้งกัน

ตัวอย่างในกรณีที่สอง hallucination คำตอบของคำถามเดียวกัน ที่ถาม LLMs ตั้งแต่สอง ครั้งขึ้นไป และให้ response ของทั้งสองครั้งที่มีความขัดแย้งกัน จากการทดลอง พบกรณีนี้กับ Perplexity, Gemini และ Copilot สำหรับกรณีคำถามเกี่ยวกับ The Temple of Dawn คำตอบจาก Perplexity ในครั้งแรกระบุไม่สามารถเดินทางไปโดยรถขนส่งสาธารณะ ในขณะที่ คำตอบในครั้งที่สอง ระบุวิธีการเดินทางโดยขนส่งสาธารณะโดยละเอียด ได้แก่ ทางรถไฟลอยฟ้า เมโทร เรือข้ามฟาก และรถแท็กซี่ และพบในกรณีคล้ายกันเกิดขึ้นกับ Gemini ต่อคำถามเกี่ยวกับ The Asiatique the Riverfront ซึ่งสามารถสรุปได้ว่าเกิด hallucination

อีกตัวอย่างหนึ่งเกิดกับ Gemini เกี่ยวกับคำถามเกี่ยวกับ The Temple of Dawn พบว่า รายงานจาก Gemini มีความขัดแย้งกัน เกี่ยวกับเวลาปิดทำการ ในครั้งแรก Gemini ระบุปัญหา ปิดทำการเร็วในช่วงบ่าย ส่วนคำตอบในครั้งที่สองระบุปิดทำการช่วงเย็นเวลา 18:00 น.

- 3) **ข้อบกพร่องจากข้อมูลล้าสมัย หรือเป็นความคิดเห็นส่วนบุคคล และข้อมูลที่มีจำกัด** จากการทดลองพบว่า Perplexity ระบุ The Grand Palace ไม่สร้างประสบการณ์ที่แปลกใหม่ไปกว่า การเยี่ยมชมวัดอื่นในเอเชียตะวันออกเฉียงใต้ response ลักษณะนี้เป็นการใช้ความคิดเห็นส่วนตัว ของบุคคลใดบุคคลหนึ่งในการตอบคำถาม ในกรณีนี้พบอีกด้วยว่าเอกสารที่ Perplexity ใช้อ้างอิง มาจากบล็อกส่วนตัวที่เผยแพร่ในปี 2013 ซึ่งเป็นข้อมูลที่เก่าเกินกว่าที่ควรจะนำมาใช้
- 4) **ข้อบกพร่องจากการตีความผิดโดย LLM** ในกรณีนี้พบในการทดลองว่า Copilot ให้ response เกี่ยวกับ The Grand Palace ว่าไม่เหมาะสมกับผู้พิการทางการเคลื่อนไหว ซึ่งเป็นการตีความ ข้อมูลที่ใช้อ้างอิงผิด ซึ่งในความเป็นจริงเอกสารอ้างอิงกล่าวว่า นักท่องเที่ยวสามารถขอยืมรถเข็น (wheelchair) ให้กับผู้พิการ หากจำเป็นต้องใช้ได้

สำหรับงานวิจัยในอนาคต ข้อมูลจาก RAG และ LLMs ในส่วนที่ถูกพิสูจน์ความน่าเชื่อถือแล้วจะถูกนำมาเสนอเป็นแนวคิดในส่วนนี้ เพื่อเสนอต่อรัฐบาลให้นำไปใช้ประโยชน์ต่อเป็นการพัฒนาท้องถิ่น รวมถึงออกกฎหรือแนวปฏิบัติตามความเหมาะสม

ตารางที่ 6.1 แสดงแนวคิดในการพัฒนาสำหรับสถานที่ท่องเที่ยวทั้ง 6 แห่งในกรุงเทพฯ

สถานที่ท่องเที่ยว	แนวทางการพัฒนา
The Grand Palace	- ปรับปรุงราคาบัตรเข้าชมให้มีตัวเลือกที่หลากหลาย - พัฒนาข้อมูลประกอบการท่องเที่ยว - ฝึกอบรมเจ้าหน้าที่ให้มีความเข้าใจถึงความหลากหลายทางเชื้อชาติและวัฒนธรรม - ผ่อนคลายมาตรการการแต่งกาย
The Temple of the Reclining Buddha	เพิ่มมาตรการเข้มงวดด้านความปลอดภัย
The Temple of Dawn	แนะนำช่วงเวลาที่เหมาะสมในการท่องเที่ยว
Khaosan Road	- ปรับปรุงพื้นที่เอกลักษณ์ความเป็นไทย - วางมาตรการและกรอบการค้าและการเสกฏาเพื่อการสันทานการ
Chatuchak Weekend Market	- ทำป้าย แผนที่ หรือแอปพลิเคชัน ที่ช่วยให้นักท่องเที่ยวเดินสำรวจภายในสถานที่ได้สะดวกสบาย - สร้างใบรับรองมาตรฐานสินค้าที่เชื่อถือได้
Asiatique the Riverfront	- ฟื้นฟูเอกลักษณ์ของตลาดแบบไทยแท้ เพื่อสร้างประสบการณ์ความเป็นไทยให้นักท่องเที่ยว - สนับสนุนให้มีการแข่งขันราคาของร้านค้าและร้านอาหารเพื่อประโยชน์ผู้บริโภค

ตารางที่ 6.1 เป็นข้อเสนอแนะสำหรับปรับปรุงสถานที่ท่องเที่ยวทั้ง 6 แห่งในกรุงเทพฯ จากข้อมูลที่ได้จากการประมวลผลโดย RAG และ LLMs ซึ่งเป็นข้อมูลที่ได้ผ่านการตรวจสอบความน่าเชื่อถือได้

ในส่วนของ The Grand Palace ควรมีการปรับค่าเข้าชม ซึ่งอาจทำให้ถูกลงได้ในกรณีที่บางส่วนของอาคาร หรือสถานที่บางแห่งภายในปิดปรับปรุงหรือปิดเพื่อเตรียมงานพระราชพิธีสำคัญ และตามที่นักท่องเที่ยวส่วนหนึ่งนำเสนอ ค่าเข้าชมควรทำเป็นแบบรวมในราคาเหมาจ่ายที่ทำให้สามารถเข้าเยี่ยมชมหลายสถานที่ ได้แก่ The Grand Palace, The Temple of the Reclining Buddha และ The Temple of Dawn

ด้านการอำนวยความสะดวกในการท่องเที่ยวในบริเวณ The Grand Palace ควรเตรียมข้อมูลให้กับนักท่องเที่ยวทั้งภาษาไทย และภาษาอังกฤษ ในหลากหลายรูปแบบ เช่นป้ายบอกทาง และไกด์นำเที่ยวแบบเสียง (audio tour guide) ที่ทำให้นักท่องเที่ยวเดินเที่ยวชมตามจุดต่าง ๆ และสร้างความเข้าใจระหว่างเที่ยวชมได้ด้วยตนเอง

สำหรับเจ้าหน้าที่ที่คอยให้ความสะดวกใน The Grand Palace ควรได้รับการฝึกอบรมให้เข้าใจถึงธรรมเนียมและประเพณี และความหลากหลายทางเชื้อชาติที่แตกต่างกัน สิ่งเหล่านี้จะส่งผลให้การเจ้ตือนักท่องเที่ยวเกี่ยวกับพฤติกรรมในที่ศักดิ์สิทธิ์เป็นไปอย่างเหมาะสมและสุภาพ

ด้านการแต่งกาย เนื่องจากอากาศในประเทศไทย โดยเฉพาะอย่างยิ่งตอนกลางวันจะร้อนมาก ควรมีมาตรการผ่อนปรนให้นักท่องเที่ยวได้แต่งกายที่ยังรู้สึกสบายต่อสภาวะอากาศแต่ก็ยักรักษาความสุภาพเรียบร้อยไว้ได้

สำหรับ The Temple of the Reclining Buddha นักท่องเที่ยวมีความกังวลมากที่สุดเกี่ยวกับความปลอดภัย โดยเฉพาะในฤดูกาลท่องเที่ยว (high season) นักท่องเที่ยวบางคนมีประสบการณ์โดนมีจาชีพลั้งกระเป่า ในกรณีนี้ควรมีตำรวจท่องเที่ยวประจำเฉพาะสถานที่นี้ และควรเจ้ตือนักท่องเที่ยวเกี่ยวกับกลโกง เช่น ไกด์เถื่อน และนักล้งกระเป่า ผ่านคำแนะนำการท่องเที่ยว

นักท่องเที่ยวส่วนมากที่เคยเยี่ยมชม The Temple of Dawn แสดงความประทับใจอย่างรักัตาม สถานที่เที่ยวแห่งนี้ต้องการการชมวิวาทศน์กลางเจ้จ จึงควรให้คำแนะนำนักท่องเที่ยวเกี่ยวกับเวลาการท่องเที่ยวที่เหมาะสม เช่นควรแนะนำในคู่มือการท่องเที่ยวว่าให้นักท่องเที่ยวเข้าชมด้านในวัดก่อนเที่ยง และเดินเล่นรอบบริเวณวัดในช่วงเย็น เนื่องจากประเทศไทยมีลักษณะอากาศร้อนชื้น จึงควรแนะนำให้นักท่องเที่ยวหลีกเลี่ยงการมาเที่ยวชมในช่วงเที่ยงวันและช่วงบ่าย (13:00 – 15:00 น.) เพื่อหลีกเลี่ยงอากาศร้อนจัด

การท่องเที่ยวบริเวณ Khaosan Road ควรเน้นการพัฒนาธรรมไทยแท้ซึ่งเป้นความคาดหวังของนักท่องเที่ยว เนื่องจากสภาวะปัจจุบันเน้นการจัดปาร์ตี้ของสถานบันเทิงมากจนเกินไป ทั้งนี้การสร้างบรรยากาศที่แสดงเอกลักษณ์ความเป็นไทย ต้องอาศัยความร่วมมือระหว่างภาครัฐและพ่อค้าแม่ค้า เพื่อให้สินค้า ร้านค้า และแสดงโชว์มีความเด่นชัดด้านเสน่ห์ทางวัฒนธรรม อีกประการหนึ่งคือกฎระเบียบเกี่ยวกับการขายกัญชา ซึ่งควรวางกรอบการขายและการใช้อยู่ในบริเวณที่ชัดเจน และเพื่อสร้างความมั่นใจให้กับนักท่องเที่ยว จึงควรเผยแพร่กฎระเบียบการขายและการใช้กัญชาที่ชัดเจนให้กับนักท่องเที่ยวได้รับทราบ

Chatuchak Weekend Market เป็นตลาดที่มีชื่อเสียงแห่งหนึ่งของประเทศไทย อย่างไรก็ตาม นักท่องเที่ยวมักมีความสับสนในการเดินเที่ยวหาซื้อสินค้าและของที่ระลึก เนื่องจากตลาดแห่งนี้มีขนาดใหญ่มาก วิธีการแก้ปัญหา ได้แก่ การจัดทำป้ายบอกทางเป็นภาษาอังกฤษและภาษาอื่นๆ หรืออาจทำในรูปแบบแอปพลิเคชันบนโทรศัพท์เคลื่อนที่เพื่อช่วยเหลือนักท่องเที่ยวให้มีข้อมูล และแผนที่ประกอบการเดินท่องเที่ยว

ปัญหาอีกประการหนึ่งของ Chatuchak Weekend Market คือการขายของเลียนแบบในตลาดซึ่งเป็นปัญหาสำคัญเกี่ยวกับทรัพย์สินทางปัญญา เพื่อเป็นการป้องกันปัญหานี้ทางภาครัฐควรออกใบรับรองสินค้าไทยแท้ที่มีการผลิตตามมาตรฐานที่ถูกต้อง และมีเครื่องหมายสินค้าที่ชัดเจนสำหรับผู้ค้า และสินค้าที่ผ่านเกณฑ์รับรอง สามารถแสดงใบรับรองที่หน้าร้าน เพื่อเป็นข้อมูลประกอบการตัดสินใจซื้อสินค้าของนักท่องเที่ยวได้

ตลาด ร้านอาหาร และร้านค้าท้องถิ่นที่ Asiatique the Riverfront มีข้อได้เปรียบทางด้านวิวทิวทัศน์ของแม่น้ำเจ้าพระยา อย่างไรก็ตาม นักท่องเที่ยวยังคงคาดหวังที่จะเห็นเอกลักษณ์ของตลาดไทยแท้จากสถานที่ท่องเที่ยวแห่งนี้ การส่งเสริมภาพลักษณ์ของไทยจึงช่วยเพิ่มคุณค่าให้กับสถานที่และจะสามารถดึงดูดนักท่องเที่ยวต่างชาติได้มากยิ่งขึ้น ทางด้านราคาสินค้าและอาหารที่แพงกว่าปกติที่ขายในสถานที่อื่น เป็นอีกปัญหาที่ต้องมีการจัดการ ทางรัฐบาลและผู้จัดการสถานที่ควรส่งเสริมให้มีการแข่งขันราคาของร้านอาหาร และร้านค้า เพื่อประโยชน์สูงสุดของนักท่องเที่ยวให้สามารถได้ของดีในราคาที่ถูกลง

บทที่ 7

บทสรุปงานวิจัย

งานวิจัยฉบับนี้แสดงให้เห็นถึงการออกแบบ Retrieval Augmented Generation (RAG) และรายงานประสิทธิภาพความน่าเชื่อถือของผลลัพธ์ที่ได้จาก RAG ซึ่งสูงกว่า Large Language Models (LLMs) ที่มีชื่อเสียงทั้ง 4 ตัว ได้แก่ ChatGPT, Perplexity, Gemini และ Copilot ในการทดลองของงานวิจัยนี้ ยิ่งไปกว่านั้นผลการทดลองยังแสดงให้เห็นว่า เอาต์พุตจาก RAG มีความแม่นยำสูง และมีเอาต์พุตบางประการที่ใหม่ ไม่เคยปรากฏในเอาต์พุตจาก LLMs มาก่อน ในทางกลับกัน เอาต์พุตบางตัวจาก LLMs แสดงข้อผิดพลาด (error) ในลักษณะที่เรียกว่าภาพหลอน (hallucination) ที่ไม่อยู่บนพื้นฐานของข้อมูลที่แท้จริง ซึ่งปัญหานี้มาจากข้อมูลที่ใช้สอนโมเดลของ LLMs และปัญหาของการประมวลผลของโมเดลเอง ในท้ายที่สุด งานวิจัยนี้้นำเอาต์พุต ในส่วนที่มีความน่าเชื่อถือสูงจากโมเดล RAG และ LLMs มาประมวลเป็นข้อเสนอแนะสำหรับสร้างแผนพัฒนาการท่องเที่ยวที่สำคัญ และเป็นที่ยอมรับในประเทศไทย



บรรณานุกรม

- [1] Huang, K. H., Laban, P., Fabbri, A. R., Choubey, P. K., Joty, S., Xiong, C., & Wu, C. S. (2023). Embrace divergence for richer insights: A multi-document summarization benchmark and a case study on summarizing diverse information from news articles. arXiv preprint arXiv:2309.09369.
- [2] Romano, M. F., Shih, L. C., Paschalidis, I. C., Au, R., & Kolachalama, V. B. (2023). Large Language Models in Neurology Research and Future Practice. *Neurology*, 101(23), 1058–1067. <https://doi.org/10.1212/WNL.0000000000207967>
- [3] Wen, J., & Wang, W. (2023). The future of ChatGPT in academic research and publishing: A commentary for clinical and translational medicine. *Clinical and translational medicine*, 13(3), e1207. <https://doi.org/10.1002/ctm2.1207>
- [4] Tabeling, J. (2024, April 25). Perplexity AI: Exploring AI-powered search beyond Google Search Engine Land. <https://searchengineland.com/perplexity-ai-exploring-ai-powered-search-beyond-google-439879>
- [5] Team, G., Anil, R., Borgeaud, S., Wu, Y., Alayrac, J. B., Yu, J., ... & Ahn, J. (2023). Gemini: a family of highly capable multimodal models. arXiv preprint arXiv:2312.11805.
- [6] Filipsson, F. (2024, March 10). The Microsoft Copilot AI: a deep dive. Redress Compliance - Just another WordPress site. <https://redresscompliance.com/the-microsoft-copilot-ai-a-deep-dive/>
- [7] Huang, L., Yu, W., Ma, W., Zhong, W., Feng, Z., Wang, H., ... & Liu, T. (2023). A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions. arXiv preprint arXiv:2311.05232.
- [8] Mousavi, S. M., Alghisi, S., & Riccardi, G. (2024). Is Your LLM Outdated? Benchmarking LLMs & Alignment Algorithms for Time-Sensitive Knowledge. arXiv preprint arXiv:2404.08700.
- [9] Touvron, H., Martin, L., Stone, K., Albert, P., Almahairi, A., Babaei, Y., ... & Scialom, T. (2023). Llama 2: Open foundation and fine-tuned chat models. arXiv preprint arXiv:2307.09288.
- [10] Hoti, M. H., & Ajdari, J. Sentiment Analysis Using the Vader Model for Assessing Company Services Based on Posts on Social Media. *SEEU Review*, 18(2), 19-33.
- [11] Giray, L. (2023). Prompt engineering with ChatGPT: a guide for academic writers. *Annals of biomedical engineering*, 51(12), 2629-2633.
- [12] Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., ... & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in Neural Information Processing Systems*, 33, 9459-9474.
- [13] Gao, Y., Xiong, Y., Gao, X., Jia, K., Pan, J., Bi, Y., ... & Wang, H. (2023). Retrieval-augmented generation for large language models: A survey. arXiv preprint arXiv:2312.10997.
- [14] Merritt, R. (2024, April 2). What Is Retrieval-Augmented Generation aka RAG | NVIDIA Blogs. <https://blogs.nvidia.com/blog/what-is-retrieval-augmented-generation/>
- [15] Al Ghadban, Y., Lu, H. Y., Adavi, U., Sharma, A., Gara, S., Das, N., ... & Hirst, J. E. (2023). Transforming healthcare education: Harnessing large language models for frontline health worker capacity building using retrieval-augmented generation. *medRxiv*, 2023-12.

- [16] Zhang, B., Yang, H., Zhou, T., Ali Babar, M., & Liu, X. Y. (2023, November). Enhancing financial sentiment analysis via retrieval augmented large language models. In Proceedings of the Fourth ACM International Conference on AI in Finance (pp. 349-356).
- [17] Araci, D. (2019). Finbert: Financial sentiment analysis with pre-trained language models. arXiv preprint arXiv:1908.10063.
- [18] Dong, C. (2023). How to Build an AI Tutor that Can Adapt to Any Course and Provide Accurate Answers Using Large Language Model and Retrieval-Augmented Generation. arXiv preprint arXiv:2311.17696.
- [19] Kulkarni, M., Tangarajan, P., Kim, K., & Trivedi, A. (2024). Reinforcement learning for optimizing rag for domain chatbots. arXiv preprint arXiv:2401.06800.



ไม่มีเนื้อหาจากต้นฉบับ



ประวัติผู้วิจัย



1. ชื่อ-สกุล นางสาวพรภัทร์ ศิริธรรมกุล
2. ตำแหน่งทางวิชาการ อาจารย์
3. ตำแหน่งปัจจุบัน อาจารย์ประจำสาขาวิชาวิศวกรรมคอมพิวเตอร์
คณะวิศวกรรมศาสตร์
มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร

4. วุฒิการศึกษา

ระดับการศึกษา	ชื่อวุฒิ	วิชาเอก	สถาบันการศึกษา	ปีที่สำเร็จการศึกษา
ปริญญาตรี	วศ.บ.	วิศวกรรมคอมพิวเตอร์	สถาบันเทคโนโลยีพระจอมเกล้าเจ้าคุณทหารลาดกระบัง	2548
ปริญญาโท	วท.ม.	วิศวกรรมซอฟต์แวร์	จุฬาลงกรณ์มหาวิทยาลัย	2552
ปริญญาโท	M.S.	Information Systems and Technology	Claremont Graduate University, USA	2557
ปริญญาเอก	Ph.D.	Information Systems and Technology	Claremont Graduate University, USA	2559

5. สถานที่ติดต่อ

ที่ทำงาน

สาขาวิชาวิศวกรรมคอมพิวเตอร์ คณะวิศวกรรมศาสตร์

มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร โทรศัพท์ (02) 02-836-3000 ต่อ 4183 หรือ 4184

อีเมล pornpat.s@rmutp.ac.th และ sirithumgul.p@gmail.com

ที่อยู่ 1381 ถนนประชาราษฎร์ 1 แขวงวงศ์สว่าง เขตบางซื่อ กทม. 10800

6. ความชำนาญ และความสนใจพิเศษ Software Engineering, Data Science, Natural Language Processing และ Artificial Intelligence

7. พทุนการศึกษา

- พทุนรัฐบาลไทยสังกัดกระทรวงวิทยาศาสตร์และเทคโนโลยี (ปี 2552 - 2559) ประเภท บุคคลภายนอก เพื่อศึกษาต่อในระดับปริญญาโท – เอก ณ ประเทศสหรัฐอเมริกา

- **ทุนนักวิจัยใหม่กระทรวงวิทยาศาสตร์ (วท.) (ปี 2560)** สำนักงานพัฒนาวิทยาศาสตร์และเทคโนโลยีแห่งชาติ (สวทช.) หัวข้อวิจัย “A Model for Constructing an Expert System Used for Generating Gap-fill Multiple Choice Questions”
- **ทุนนักวิจัยแลกเปลี่ยนภายใต้ความร่วมมือระหว่างสำนักงานการวิจัยแห่งชาติ (วช.) ประเทศไทยและ Indian Council of Social Science Research (ICSSR) ประเทศอินเดีย (ปี 2564)** หัวข้อวิจัย “A Study of Applying a Knowledge Expert System to Enhance Education Focusing on Information Technology and Software Engineering in Thailand and India”

8. ผลงานทางวิชาการ

8.1 งานประพันธ์ตำราเรียน

- พรภัทร์ ศิริธรรมกุล (2565).** หน่วยที่ 8. ปัญญาประดิษฐ์สำหรับการจัดการโลจิสติกส์และห่วงโซ่อุปทาน. ใน *ตำราเรียนชุดวิชาปัญญาประดิษฐ์สำหรับอุตสาหกรรมและธุรกิจ หน่วยที่ 8-15*. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมมาธิราช.
- พรภัทร์ ศิริธรรมกุล (2565).** หน่วยที่ 11. ปัญญาประดิษฐ์สำหรับสร้างประสิทธิภาพให้การจัดการกระบวนการทางธุรกิจ. ใน *ตำราเรียนชุดวิชาปัญญาประดิษฐ์สำหรับอุตสาหกรรมและธุรกิจ หน่วยที่ 8-15*. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมมาธิราช.
- พรภัทร์ ศิริธรรมกุล (2565).** หน่วยที่ 15. ปัญญาประดิษฐ์สำหรับการวิเคราะห์ผลป้อนกลับจากลูกค้า. ใน *ตำราเรียนชุดวิชาปัญญาประดิษฐ์สำหรับอุตสาหกรรมและธุรกิจ หน่วยที่ 8-15*. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมมาธิราช.
- พรภัทร์ ศิริธรรมกุล และภรณ์ ศรีสุทธิ (2563).** หน่วยที่ 7. วิธีการเพื่อการพัฒนากระบวนทัศน์. ใน *ตำราเรียนชุดวิชาเทคโนโลยีเพื่อการจัดการสารสนเทศสำหรับนักศึกษาระดับปริญญาโท หน่วยที่ 6-10*. นนทบุรี: สาขาวิชาศิลปศาสตร์ แขนงวิชาสารสนเทศศาสตร์ สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมมาธิราช.
- พรภัทร์ ศิริธรรมกุล และชุตติ์ ธนภักดิ์ภาดา (2563).** หน่วยที่ 11. การจัดหา การติดตั้ง การบำรุงรักษา และการประเมินระบบสารสนเทศ. ใน *ตำราเรียนชุดวิชาเทคโนโลยีสารสนเทศเบื้องต้น หน่วยที่ 11-15*. นนทบุรี: สาขาวิชาศิลปศาสตร์ แขนงวิชาสารสนเทศศาสตร์ สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมมาธิราช.

พรภัทร์ ศิริธรรมกุล (2563). หน่วยที่ 3. ทฤษฎีการรับรู้. ใน *ตำราเรียนชุดวิชาการออกแบบส่วนปฏิสัมพันธ์บนเว็บและโมบาย* หน่วยที่ 1-7. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมาธิราช.

พรภัทร์ ศิริธรรมกุล (2563). หน่วยที่ 8. การออกแบบประสบการณ์ผู้ใช้. ใน *ตำราเรียนชุดวิชาการออกแบบส่วนปฏิสัมพันธ์บนเว็บและโมบาย* หน่วยที่ 8-15. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมาธิราช.

พรภัทร์ ศิริธรรมกุล (2563). หน่วยที่ 13. เซอร์วิสบนแอนดรอยด์. ใน *ตำราเรียนชุดวิชาการโปรแกรมประยุกต์บนอุปกรณ์เคลื่อนที่* หน่วยที่ 11-15. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมาธิราช.

พรภัทร์ ศิริธรรมกุล. (2562). หน่วยที่ 2. ระดับและกระบวนการด้านความมั่นคงปลอดภัยไซเบอร์. ใน *ตำราเรียนชุดวิชาความมั่นคงปลอดภัยไซเบอร์* หน่วยที่ 1-5. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมาธิราช.

พรภัทร์ ศิริธรรมกุล (2562). หน่วยที่ 4. การเก็บรวบรวมและการเตรียมข้อมูลก่อนการประมวลผล. ใน *ตำราเรียนชุดวิชาการวิเคราะห์ข้อมูล* หน่วยที่ 1-5. นนทบุรี: สาขาวิชาวิทยาศาสตร์และเทคโนโลยี สำนักพิมพ์มหาวิทยาลัยสุโขทัยธรรมาธิราช.

8.2 ผลงานวิจัยที่เผยแพร่

Sirithumgul, P. (2023). Unlocking the Potential of ChatGPT: A Grounded Theory Exploration of its Impact on the Business Landscape. In *Proceedings of the 29th Americas Conference on Information Systems*, Panama City, Panama. August 10 – 12, 2023.

Sirithumgul, Pornpat and Prasertsilp, Pimpaka. (2022). Applying Sentiment Analysis and Machine Learning Algorithms on Students' Reflections to Identify an Effective Teaching Strategy as a Factor of Learning Successes. In *Proceedings of the 2022 Australasian Conference on Information Systems (ACIS 2022)*. <https://aisel.aisnet.org/acis2022/67>.

Sirithumgul, P., Prasertsilp, P., & Olfman, L. (2022). An Algorithm for Generating Gap-Fill Multiple Choice Questions of an Expert System. In *Proceedings of the 55th Hawaii International Conference on System Sciences*, Hawaii, USA. January 4-7, 2022

Sirithumgul, P., Prasertsilp, P., & Olfman, L. (2021). An Add-On for Empowering Google Forms to be an Automatic Question Generator in Online Assessments. arXiv preprint arXiv:2110.15220.

Sirithumgul, Pornpat, Prasertsilp, Pimpaka, Suksa-Ngiam, Watanyoo & Olfman, Lorne. An Ontology-Based Framework as a Foundation of an Information System for Generating Multiple-Choice Questions. In Proceedings of the 25th Americas Conference on Information Systems, Cancun, Mexico. August 15 – 17, 2019.

Sirithumgul, P. (2016). An ontology-based algorithm as a foundation of an automated knowledge assessment tool applied in the scientific discussions. Available from ProQuest Dissertations & Theses Global.

Sirithumgul, Pornpat, & Olfman, Lorne. A model for measuring knowledge constructions of students in online discussions. In Proceedings of the 24th Australasian Conference on Information Systems, Melbourne, Australia. December 4-6, 2013.

Sirithumgul, Pornpat, Suchato, Atiwong, & Punyabukkana, Proadpran. Quantitative Evaluation for Web Accessibility with Respect to Disabled Groups. In Proceedings of the 6th International Cross- Disciplinary Conference on Web Accessibility (W4A2009), Madrid, Spain. April 20-21, 2009.

Sirithumgul, Pornpat, Suchato, Atiwong, & Punyabukkana, Proadpran. Elearning Transformation for the visually impaired. In Proceedings of the 8th Distance Learning and the Internet Conference (DLI2007), Bangkok, Thailand. December 12-15, 2007.

8.3 โครงการวิจัย

พรภัทร์ ศิริธรรมกุล (2566). หัวหน้าโครงการวิจัยเรื่อง *บทบาทของบทวิพากษ์วิจารณ์จากลูกค้าบนสื่อสังคมต่อการสนับสนุนธุรกิจท่องเที่ยวในประเทศไทย: เทคนิคการวิเคราะห์ข้อมูลขนาดใหญ่* (A Role of Customer Reviews on Social Media to Promoting Travel Businesses in Thailand: A Technique for Analyzing Big Data) สนับสนุนโดยทุน

งบประมาณเงินรายได้หน่วยงาน มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร ประจำปี 2566

พรภัทร์ ศิริธรรมกุล (2565). หัวหน้าโครงการวิจัยเรื่อง การประยุกต์อัลกอริทึมการเรียนรู้ของเครื่องเพื่อระบุปัจจัยความสำเร็จที่ถูกสนับสนุนโดยการเรียนออนไลน์แบบวนซ้ำ (An Application of Machine Learning Algorithms to Identify Success Factors being Promoted by Repetition Online Learning) สนับสนุนโดยทุนงบประมาณเงินรายได้หน่วยงาน มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร ประจำปี 2565

พรภัทร์ ศิริธรรมกุล (2563). หัวหน้าโครงการวิจัยเรื่อง กรอบงานบนพื้นฐานของภววิทยาในฐานะที่เป็นพื้นฐานของระบบสารสนเทศสำหรับสร้างคำถามแบบปรนัย (An Ontology-based Framework as a Foundation of an Information System for Generating Multiple-choice Questions) สนับสนุนโดยทุนงบประมาณเงินรายได้หน่วยงาน มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร ประจำปี 2563

พรภัทร์ ศิริธรรมกุล (2562). หัวหน้าโครงการวิจัยเรื่อง แบบจำลองสำหรับสร้างระบบผู้เชี่ยวชาญที่ใช้สำหรับผลิตคำถามแบบเติมคำในช่องว่างที่มีตัวเลือกแบบปรนัย (A Model for Constructing an Expert System Used for Generating Gap-fill Multiple Choice Questions) สนับสนุนโดยโครงการสนับสนุนทุนนักวิจัยใหม่ กระทรวงวิทยาศาสตร์ (วท.) สำนักงานพัฒนาวิทยาศาสตร์และเทคโนโลยีแห่งชาติ (สวทช.)

8.4 ประสพการณ์การเป็นวิทยากรและอาจารย์พิเศษ

พรภัทร์ ศิริธรรมกุล (2563). วิทยากรบรรยายพิเศษเรื่อง การบริการเว็บและเอพีไอบนเว็บ (Web Services and Web APIs) ซึ่งเป็นส่วนหนึ่งของชุดวิชา การโปรแกรมเว็บ (Web Programming) สำหรับนักศึกษาระดับปริญญาตรี

พรภัทร์ ศิริธรรมกุล (2562). วิทยากรรายการโทรทัศน์เพื่อการสอนทางไกลเรื่อง ระดับและกระบวนการด้าน ความมั่นคงปลอดภัยไซเบอร์ ซึ่งเป็นส่วนหนึ่งของชุดวิชา ความมั่นคงปลอดภัยไซเบอร์ แพร่ภาพผ่าน สถานีวิทยุโทรทัศน์ของมหาวิทยาลัยสุโขทัยธรรมาธิราช (STOU Channel) สถานีวิทยุโทรทัศน์ การศึกษาทางไกลผ่านดาวเทียม True Vision ช่อง DLTV14 ทางเครือข่ายอินเทอร์เน็ต www.stou.ac.th และสื่อสังคม YouTube

พรภัทร์ ศิริธรรมกุล (2562). วิทยากรรายการโทรทัศน์เพื่อการสอนทางไกลเรื่อง การเก็บรวบรวมและการเตรียมข้อมูลก่อนการประมวลผล ซึ่งเป็นส่วนหนึ่งของชุดวิชา การวิเคราะห์ข้อมูล แพร่ภาพผ่าน สถานีวิทยุโทรทัศน์ของมหาวิทยาลัยสุโขทัยธรรมาธิราช (STOU Channel)

สถานีวิทยุโทรทัศน์ การศึกษาทางไกลผ่านดาวเทียม True Vision ช่อง DLTV14 ทาง
เครือข่ายอินเทอร์เน็ต www.stou.ac.th และสื่อสังคม YouTube

พรภัทร์ ศิริธรรมกุล (2562). อาจารย์อาสาและอาจารย์พิเศษวิชา *การโปรแกรมคอมพิวเตอร์เบื้องต้น โดยใช้ภาษาซี* ให้กับโครงการการศึกษาทางไกลผ่านดาวเทียมไกลกังวล ภายใต้การสนับสนุนของมูลนิธิการศึกษาทางไกลผ่านดาวเทียมในพระบรมราชูปถัมภ์ ที่เปิดให้ความรู้กับนักเรียนระดับอาชีวศึกษาและประชาชนทั่วไป แพร่ภาพทางช่อง สศทท.13 (dltv13) ระหว่างวันที่ 20 มิถุนายน ถึง 13 ตุลาคม 2562

พรภัทร์ ศิริธรรมกุล (2551). วิทยากรอบรมหัวข้อ Web Accessibility, WCAG 2.0 and Assistive Technology ในโครงการความร่วมมือระหว่าง คณะวิศวกรรมศาสตร์ จุฬาลงกรณ์มหาวิทยาลัยและกระทรวงเทคโนโลยีสารสนเทศและการสื่อสาร

พรภัทร์ ศิริธรรมกุล (2551). วิทยากรอบรมการใช้งานโปรแกรมอ่านหน้าจอคอมพิวเตอร์ Non-Visual Desktop Access (NVDA) ให้กับนักเรียนตาบอด ณ ศูนย์การศึกษาพิเศษส่วนกลาง กรุงเทพมหานคร

8.5 ประสบการณ์การสอน

1. อาจารย์ผู้สอนวิชา *วิทยาศาสตร์ข้อมูล (Data Science)* สำหรับนักศึกษาชั้นปีที่ 4 สาขาวิชาวิศวกรรมคอมพิวเตอร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร (ตั้งแต่ปี 2562 - ปัจจุบัน)
2. อาจารย์ผู้สอนวิชา *พื้นฐานการออกแบบและการโปรแกรมเชิงวัตถุโดยใช้ภาษาไพธอน (Fundamentals of Object-Oriented Design and Programming Using Python)* สำหรับนักศึกษาชั้นปีที่ 2 ภาควิชาวิศวกรรมคอมพิวเตอร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร (ตั้งแต่ปี 2562 - ปัจจุบัน)
3. อาจารย์ผู้สอนวิชา *วิศวกรรมซอฟต์แวร์ (Software Engineering)* สำหรับนักศึกษาชั้นปีที่ 3 ภาควิชาวิศวกรรมคอมพิวเตอร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร (ตั้งแต่ปี 2559 - ปัจจุบัน)
4. อาจารย์ผู้สอนวิชาการโปรแกรมคอมพิวเตอร์ สำหรับนักศึกษาชั้นปีที่ 1 ภาควิชาวิศวกรรมคอมพิวเตอร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร (ตั้งแต่ปี 2559 - 2562)
5. ผู้ช่วยสอนวิชา Knowledge Management and Social Media (Transdisciplinary Course) สำหรับนักศึกษาระดับปริญญาเอก Claremont Graduate University (ภาคการศึกษา Fall 2014)
6. ผู้ช่วยสอนวิชาวิศวกรรมซอฟต์แวร์ สำหรับนักศึกษาชั้นปีที่ 3 ภาควิชาวิศวกรรมคอมพิวเตอร์ จุฬาลงกรณ์มหาวิทยาลัย (ปีการศึกษา 2550 - 2552)

8.6 งานบริการวิชาการ

1. **Program committee member and reviewer:** The International Multi-Conference on Computing in the Global Information Technology (ICCGI)
2. **Session chair:** The 17th International Conference on Electrical Engineering/Electronics, Computer, Telecommunications and Information Technology (ECTI-CON 2020)
3. **Reviewer:**
 - 3.1 Computer & Education: An International Journal
 - 3.2 International Conference on Information Systems (ICIS)
 - 3.3 Hawaii International Conference on System Sciences (HICSS)
 - 3.4 Americas Conference on Information Systems (AMCIS)

8.7 Professional Society Membership

1. Association for Computing Machinery (ACM) (2009 – present)
2. Association for Information Systems (AIS) (2019 – present)

9. ประสบการณ์การทำงานภาคเอกชน

1. **ตำแหน่ง Senior software engineer** บริษัท Motif Technology Co., Ltd. เพื่อสร้างระบบการจัดการสารสนเทศ (DIS-MIS) ให้กับกรมสอบสวนคดีพิเศษแห่งราชอาณาจักรไทย (Department of Special Investigation: DSI)
2. **ตำแหน่ง Software engineer** บริษัท Motif Technology Co., Ltd. เพื่อสร้างระบบการจัดการบัญชีหลัก (Mainstream Account System: MAS) ให้กับบริษัท Bangkok Aviation Fuel Services Public Company Limited (BAFS)

10. อื่น ๆ

เป็นผู้แทนคณะวิศวกรรมศาสตร์ มหาวิทยาลัยเทคโนโลยีราชมงคลพระนคร ทำหน้าที่ดูแลและฝึกสอนนักศึกษาแลกเปลี่ยนโครงการ International Association for the Exchange of Students for Technical Experience (IAESTE)