

Journal of Engineering and Digital Technology (JEDT)

Thai-Nichi Institute of Technology

Vol.13 No.2

July - December 2025



วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล

สถาบันเทคโนโลยีไทย-ญี่ปุ่น

ปีที่ 13 ฉบับที่ 2 กรกฎาคม - ธันวาคม 2568

วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล

Journal of Engineering and Digital Technology (JEDT)

ปีที่ 13 ฉบับที่ 2 กรกฎาคม – ธันวาคม 2568 Vol. 13 No. 2 July - December 2025

ISSN (Online) 2774-0617

ความเป็นมา

ด้วยสถาบันเทคโนโลยีไทย-ญี่ปุ่น มีนโยบายสนับสนุนการเผยแพร่บทความจากผลงานวิจัยที่มีคุณภาพ เพื่อเป็นประโยชน์ในการพัฒนาความรู้แก่สังคม โดยเฉพาะภาคธุรกิจและอุตสาหกรรม จึงได้จัดทำวารสารวิชาการ คือ วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล (เดิมชื่อ วารสารสถาบันเทคโนโลยีไทย-ญี่ปุ่น : วิศวกรรมศาสตร์และเทคโนโลยี)

วัตถุประสงค์

1. เพื่อส่งเสริมและเผยแพร่ผลงานวิจัยสาขาวิศวกรรมศาสตร์และเทคโนโลยีที่มีคุณภาพ
2. เพื่อเป็นสื่อกลางในการแลกเปลี่ยนองค์ความรู้ทางการวิจัยในสาขาวิศวกรรมศาสตร์และเทคโนโลยี
3. เพื่อพัฒนาศักยภาพทางการวิจัยสาขาวิศวกรรมศาสตร์และเทคโนโลยี

ที่ปรึกษา

รองศาสตราจารย์รังสรรค์ เลิศในสัตย์

ขอบเขตเนื้อหา

บทความวิจัยทางด้านวิทยาศาสตร์ วิศวกรรมศาสตร์ และเทคโนโลยี ได้แก่ เทคโนโลยีวิศวกรรม เทคโนโลยีอุตสาหกรรม เทคโนโลยีผลิตภัณฑ์ เทคโนโลยีสารสนเทศ วิทยาศาสตร์ประยุกต์ วิทยาศาสตร์กายภาพ วิทยาศาสตร์ชีวภาพ วิทยาการคอมพิวเตอร์ วิทยาศาสตร์เคมี และสาขาอื่น ๆ ที่เกี่ยวข้อง

กำหนดออกเผยแพร่

วารสารตีพิมพ์เผยแพร่ราย 6 เดือน (ปีละ 2 ฉบับ)

- ฉบับที่ 1 มกราคม – มิถุนายน
- ฉบับที่ 2 กรกฎาคม – ธันวาคม

เจ้าของ

สถาบันเทคโนโลยีไทย-ญี่ปุ่น (Thai-Nichi Institute of Technology : TNI)

1771/1 ถนนพัฒนาการ แขวงสวนหลวง เขตสวนหลวง กรุงเทพฯ 10250

1771/1 Pattanakarn Rd., Suanluang, Bangkok 10250

บรรณาธิการ

อาจารย์ ดร.ก้องกิจ วรพุทธพร

สถาบันเทคโนโลยีไทย-ญี่ปุ่น

กองบรรณาธิการ

- | | |
|--|--|
| 1. ศาสตราจารย์ ดร.มานะ ศรียุทธศักดิ์ | จุฬาลงกรณ์มหาวิทยาลัย |
| 2. ศาสตราจารย์ ดร.ไพโรจน์ สิงห์นาคกิจ | จุฬาลงกรณ์มหาวิทยาลัย |
| 3. รองศาสตราจารย์ ดร.ชาญชัย ปลื้มปิติวิริยะเวช | จุฬาลงกรณ์มหาวิทยาลัย |
| 4. ผู้ช่วยศาสตราจารย์ ดร.อัศวิน ศิริสุข | จุฬาลงกรณ์มหาวิทยาลัย |
| 5. ผู้ช่วยศาสตราจารย์ สุวันชัย พงษ์สุกิจวัฒน์ | จุฬาลงกรณ์มหาวิทยาลัย |
| 6. รองศาสตราจารย์ ดร.จิรเดช เสี่ยงมศักดิ์ | มหาวิทยาลัยขอนแก่น |
| 7. ศาสตราจารย์ ดร.อิสสระชัย งามหิรัญ | สถาบันเทคโนโลยีพระจอมเกล้าเจ้าคุณทหารลาดกระบัง |
| 8. ศาสตราจารย์ ดร.วรงค์ ตั้งศรีรัตน์ | สถาบันเทคโนโลยีพระจอมเกล้าเจ้าคุณทหารลาดกระบัง |
| 9. รองศาสตราจารย์ ดร.กันต์พงษ์ วรรัตน์ปัญญา | สถาบันเทคโนโลยีพระจอมเกล้าเจ้าคุณทหารลาดกระบัง |
| 10. ศาสตราจารย์ ดร.ประยุทธ์ อัครเอกผาลิน | มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ |
| 11. รองศาสตราจารย์ ดร.พยุ่ง มีสัจ | มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าพระนครเหนือ |
| 12. ผู้ช่วยศาสตราจารย์ ดร. ศิรดา ศิริธร | มหาวิทยาลัยมหิดล |
| 13. รองศาสตราจารย์ ดร.ปริทรรศน์ พันธุบรรยงก์ | นักวิชาการอิสระ |
| 14. รองศาสตราจารย์ อัญชลี สุพิทักษ์ | สถาบันเทคโนโลยีไทย-ญี่ปุ่น |
| 15. ผู้ช่วยศาสตราจารย์ ดร.วิภาวดี วงษ์สุวรรณ | สถาบันเทคโนโลยีไทย-ญี่ปุ่น |

คณะกรรมการดำเนินงาน

- | | |
|---|-------------------------------------|
| 1. ผู้ช่วยศาสตราจารย์ ดร.รชตะ รุ่งตระกูลชัย | 7. อาจารย์ ดร.ภาสกร อภิรักษ์วรพินิต |
| 2. รองศาสตราจารย์ ดร.วรากร ศรีเชวงทรัพย์ | 8. อาจารย์วุฒิพงษ์ ปะวะสาร |
| 3. ผู้ช่วยศาสตราจารย์ ดร.ธัญมัย เจียรกุล | 9. นางสาวสุพิศ บายคายคม |
| 4. ผู้ช่วยศาสตราจารย์ อมราวดี ทัพพูน | 10. นางสาวไกล่รุ่ง หมายชื่น |
| 5. อาจารย์ ดร.อภิษฎา นัมคุ้มภัย | 11. นางสาวพิมพ์รต พิพัฒน์กุล |
| 6. อาจารย์ ดร.เอกอุ ธรรมกรบัญญัติ | 12. นางสาวอารีสา จิระเวชถาวร |

การติดต่อกองบรรณาธิการ

1771/1 ถนนพัฒนาการ แขวงสวนหลวง เขตสวนหลวง กรุงเทพฯ 10250

Tel. 0-2763-2600 (Ext. 2704, 2752)

Website : <https://ph01.tci-thaijo.org/index.php/TNJJournal>

E-mail : JEDT@tni.ac.th

บทความวิจัยที่ตีพิมพ์ในวารสารฉบับนี้เป็นความคิดเห็นส่วนตัวของผู้เขียน

กองบรรณาธิการไม่มีส่วนรับผิดชอบใด ๆ ถือเป็นความรับผิดชอบของผู้เขียนแต่เพียงผู้เดียว

บทบรรณาธิการ

กองบรรณาธิการวารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล (JEDT) ยังคงยึดมั่นต่อพันธกิจในการเป็นเวทีเผยแพร่ผลงานวิชาการที่มีคุณภาพสูง เพื่อสนับสนุนการแลกเปลี่ยนองค์ความรู้ระหว่างนักวิจัยและผู้เชี่ยวชาญในสาขาวิศวกรรมศาสตร์และเทคโนโลยีสมัยใหม่ พร้อมทั้งส่งเสริมการบูรณาการองค์ความรู้ข้ามสาขา ซึ่งให้ความสำคัญกับการเชื่อมโยงศาสตร์ที่หลากหลายเพื่อสร้างฐานความรู้ที่สามารถนำไปใช้ประโยชน์ได้จริงอย่างกว้างขวางท่ามกลางการเปลี่ยนแปลงที่รวดเร็วของเทคโนโลยีดิจิทัลในขณะนี้

ในครั้งนี้ กองบรรณาธิการขอแสดงความยินดีกับนักวิจัยทุกท่านที่ผลงานได้รับการคัดเลือกเพื่อตีพิมพ์ในวารสารฉบับนี้ ทุกบทความได้ผ่านกระบวนการกลั่นกรองและประเมินโดยผู้ทรงคุณวุฒิในแต่ละสาขาที่เกี่ยวข้องอย่างเข้มงวด พร้อมทั้งได้รับข้อชี้แนะเชิงวิชาการเพื่อยกระดับคุณภาพของผลงานวิจัยให้มีความสมบูรณ์และมีความน่าเชื่อถือในระดับสากลมากยิ่งขึ้น

ผลงานที่ได้รับการตีพิมพ์อย่างมีคุณภาพสะท้อนถึงความก้าวหน้าทางวิชาการและศักยภาพในการนำไปต่อยอดสู่การพัฒนาเทคโนโลยีและนวัตกรรมในอนาคต กองบรรณาธิการหวังเป็นอย่างยิ่งว่า นักวิจัยทุกท่านจะส่งผลงานมาให้กองบรรณาธิการได้ทำการพิจารณาเพื่อตีพิมพ์อีกในอนาคต พร้อมทั้งขอเชิญชวนนักวิจัยท่านอื่นส่งผลงานเพื่อเผยแพร่ในวารสารในโอกาสต่อไปด้วยเช่นกัน

อาจารย์ ดร.ก้องกิจ วรพุทธร
บรรณาธิการ

Editorial Message

The editorial board of the Journal of Engineering and Digital Technology (JEDT) remains committed to its mission of being a high-quality academic platform that supports the exchange of knowledge among researchers and experts in engineering and modern technology. It also promotes interdisciplinary integration, emphasizing the connection of diverse disciplines to create a knowledge base that can be widely applied in the rapid changes of digital technology today.

This time, the editorial board congratulates all researchers whose work has been selected for publication in this issue. All articles have undergone rigorous peer review and evaluation by experts in their respective fields, and have received academic guidance to further enhance the quality and credibility of their research to an international standard.

The publication of high-quality work reflects academic progress and the potential for further development of technology and innovation in the future. The editorial board sincerely hopes that all researchers will submit their work for consideration for publication again in the future, and also invite other researchers to submit their work for publication in the journal in the future as well.

Kongkij Voraputhaporn, Ph.D.
Editor-in-Chief

สารบัญ

- 1 Analysis of Foreign Tourist Review by Natural Language Processing and a Lexicon-Based Sentiment Analysis Tool
Chakkarin Santirattanaphakdi, Supakit Niwattanakul
- 21 Exploiting Transformer Network for Nail Diseases Classification and Recognition
Aekkarat Suksukont, Bunthida Chunngam, Ekachai Naowanich
- 29 Improving Efficiency of Ternary Tree to Support Different Quality of Service Levels
Warakorn Srichavengsup, Titichaya Thanamitsomboon, Kanticha Kittipeerachon
- 39 Improving the Efficiency of Vehicle Queuing and Product Conveyance through Simulation: A Case Study
Sureeporn Yeanyong, Aitnanat Plukfung, Boonsiwatt Khuamthap, Anot Chaimanee
- 55 Optimal Allocation and Deployment of Roadside Units in Cloud-Based Internet of Vehicles Framework
Nay Myo Sandar, Surekha Lanka, Thinzar Aung Win, Shuvra Tripura
- 63 Optimal 2DOF-PID Controller Design Using Whale Optimization Algorithm
Kittisak Lurang, Thiwa Jitwang, Deacha Puangdownreong
- 75 Rainfall Impact on Soil Behavior: Landslide Risk Assessment Using Physical Modeling
Laddawon Dul, Taweechai Yeemao, Suphakrit Kantakam
- 89 Solving Factory Maintenance Problem: A Case Study of a Semi-Finished Food Product Manufacturing and Distribution Company in Uttaradit Province
Adun Phuk-in

102 Transforming Unstructured Data in IT Project: A Comparative Study of Zero-Shot and Generative
AI Text Classification

Cai Tung-lersloy, Worapat Paireekreng, Nantika Prinyapol

การวิเคราะห์บทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติด้วยการประมวลผล ภาษาธรรมชาติ และเครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมเป็นฐาน

จักรินทร์ สันติรัตนภักดี^{1*} ศุภกฤษฎี นวัตกรรมกุล²

^{1*}สาขาวิชาเทคโนโลยีธุรกิจดิจิทัล คณะบริหารธุรกิจ มหาวิทยาลัยวงษ์ชวลิตกุล, นครราชสีมา, ประเทศไทย

²สำนักวิชาศาสตร์และศิลป์ดิจิทัล มหาวิทยาลัยเทคโนโลยีสุรนารี, นครราชสีมา, ประเทศไทย

*ผู้ประพันธ์บทความ อีเมล : chakkarin_san@vu.ac.th

รับต้นฉบับ : 30 มกราคม 2568; รับบทความฉบับแก้ไข : 14 มีนาคม 2568; ตอบรับบทความ : 31 มีนาคม 2568

เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

พฤติกรรมของนักท่องเที่ยวในปัจจุบันคือการค้นหาข้อมูลออนไลน์ประกอบการตัดสินใจท่องเที่ยว แม้บทวิจารณ์ออนไลน์จะมีข้อดีหลายประการ แต่การวิเคราะห์ข้อมูลจำนวนมากนั้นย่อมสิ้นเปลืองเวลาและทรัพยากรในการแยกข้อมูลที่สำคัญ ดังนั้นการวิเคราะห์ข้อความอย่างเป็นระบบจึงช่วยให้สามารถนำเสนอสารสนเทศได้ครอบคลุมและตรงประเด็นมากขึ้น งานวิจัยนี้มีวัตถุประสงค์เพื่อวิเคราะห์บทวิจารณ์ด้วยเครื่องมือวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรม และเพื่อสร้างรายงานเชิงวิเคราะห์บทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติต่ออุทยานแห่งชาติเขาใหญ่ โดยรวบรวมบทวิจารณ์เกี่ยวกับอุทยานแห่งชาติเขาใหญ่จากแพลตฟอร์มออนไลน์ รวมทั้งสิ้น 12,035 รายการ โดยใช้ TextBlob, Flair และ VADER จำแนกความคิดเห็นออกเป็นเชิงบวก เชิงลบ และเป็นกลาง ผลการทดสอบพบว่า VADER มีค่าความถูกต้องเฉลี่ยสูงที่สุดที่ 76% อย่างไรก็ตาม จากการวิเคราะห์ความรู้สึก ยังพบบทวิจารณ์ที่คะแนนความพึงพอใจขัดแย้งกับเนื้อหา ซึ่งแสดงถึงข้อจำกัดของการใช้ผลการวิเคราะห์ความรู้สึกเพียงอย่างเดียว ในการสะท้อนความคิดเห็น เนื่องจากความคิดเห็นมีความซับซ้อน และยากที่จะเปรียบเทียบความแตกต่างด้วยมาตรฐานเดียวกัน เพื่อแก้ไขปัญหาดังกล่าวจึงสร้างรายงานเชิงวิเคราะห์จากเปรียบเทียบความสัมพันธ์ระหว่างคำสำคัญและบทวิจารณ์ที่เกี่ยวข้องด้วยการวัดความคล้ายคลึงโคไซน์ และสรุปข้อความด้วยการประมวลผลภาษาธรรมชาติร่วมกับการนำเสนอภาพข้อมูล เพื่อสะท้อนถึงจุดเด่นและข้อจำกัดของอุทยาน ผลการวิเคราะห์ข้อมูล พบว่าแม้ว่าอุทยานแห่งชาติเขาใหญ่จะได้รับการชื่นชมในด้านความงามทางธรรมชาติ ความอุดมสมบูรณ์ของทรัพยากร และความหลากหลายทางชีวภาพ แต่ก็ยังมีข้อจำกัดบางประการที่ถูกกล่าวถึงในบทวิจารณ์ ได้แก่ ค่าธรรมเนียมที่แตกต่างกันระหว่างนักท่องเที่ยวชาวไทยและชาวต่างชาติ นอกจากนี้ ระบบขนส่งสาธารณะยังขาดความสะดวกและไม่เพียงพอ รวมถึงต้องการประสบการณ์การท่องเที่ยวที่คุ้มค่าและกิจกรรมที่เหมาะสมกับนักท่องเที่ยวทุกกลุ่ม ผลลัพธ์จากงานวิจัยยังเป็นหนึ่งในข้อมูลพื้นฐานในการตัดสินใจของนักท่องเที่ยว อีกทั้งเป็นประโยชน์ต่อการกำหนดแนวทาง ในการพัฒนาอุทยานแห่งชาติเขาใหญ่ให้ตอบสนองความต้องการของนักท่องเที่ยว นอกจากนี้ ยังเป็นต้นแบบในการประยุกต์ใช้กับแหล่งท่องเที่ยวอื่น ๆ เพื่อยกระดับคุณภาพการให้บริการและเสริมสร้างศักยภาพในการแข่งขันของอุตสาหกรรมการท่องเที่ยวของประเทศไทยในระยะยาว

คำสำคัญ : เครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมเป็นฐาน การประมวลผลภาษาธรรมชาติ การวิเคราะห์ความรู้สึก การวิเคราะห์ข้อความ



Analysis of Foreign Tourist Review by Natural Language Processing and a Lexicon-Based Sentiment Analysis Tool

Chakkarin Santirattanaphakdi^{1*} Suphakit Niwattanakul²

^{1*}*Digital Business Technology Program, Faculty of Business Administration, Vongchavalitkul University, Nakhon Ratchasima, Thailand*

²*Institute of Digital Arts and Science, Suranaree University of Technology, Nakhon Ratchasima, Thailand*

*Corresponding Author. E-mail address: chakkarin_san@vu.ac.th

Received: 30 January 2025; Revised: 14 March 2025; Accepted: 31 March 2025

Published online: 26 December 2025

Abstract

The current behavior of tourists is to search for online information to support travel decisions. Although online reviews have many advantages, analyzing large amounts of data is time-consuming and resource-intensive in extracting important information; therefore, systematic text analysis helps present more comprehensive and targeted information. This research aims to analyze reviews using dictionary-based sentiment analysis tools and to create an analytical report of reviews from foreign tourists toward Khao Yai National Park. A total of 12,035 reviews related to Khao Yai National Park were collected from online platforms. TextBlob, Flair, and VADER were used to classify opinions as positive, negative, or neutral. The results showed that VADER had the highest average accuracy at 76%. However, sentiment analysis also found reviews in which satisfaction scores conflicted with the content, showing the limitation of using sentiment analysis alone to reflect opinions, as opinions are complex and difficult to compare using a single standard. To address this issue, an analytical report was created by analyzing the relationships between key terms and related reviews using cosine similarity and summarizing the text using natural language processing with data visualization to reflect the strengths and limitations of the park. The analysis found that although Khao Yai National Park is praised for its natural beauty, resource richness, and biodiversity, some limitations are mentioned in the reviews, including different fees for Thai and foreign tourists, inconvenient and insufficient public transportation, and the need for valuable tourism experiences and activities suitable for all tourist groups. The results of this research provide basic information for tourist decision-making and are useful for developing guidelines to improve Khao Yai National Park to better meet tourists' needs, and can also be applied as a model for other tourist destinations to enhance service quality and the competitiveness of Thailand's tourism industry in the long term.

Keywords: Lexicon-based sentiment analysis tool, Natural language processing, Sentiment analysis, Text analysis

1) บทนำ

พัฒนาการของเทคโนโลยีสารสนเทศ และการสื่อสารส่งผลให้เกิดการเปลี่ยนแปลงพฤติกรรมการใช้ชีวิตยุคปัจจุบันไปอย่างสิ้นเชิง ไม่เว้นแม้แต่พฤติกรรมการท่องเที่ยวที่ผู้ใช้งานค้นหาข้อมูลผ่านสื่อออนไลน์ต่าง ๆ ที่นำเสนอข้อมูลอย่างหลากหลาย เพื่อสนับสนุนการท่องเที่ยว โดยเฉพาะคะแนนความพึงพอใจ (rating) ที่ได้กลายเป็นเครื่องมือสำคัญในการช่วยตัดสินใจเลือกสถานที่ท่องเที่ยว [1] จากประสบการณ์ของผู้ที่เคยเดินทางไปยังสถานที่แห่งนั้นมาแล้ว ซึ่งเพิ่มความมั่นใจในการวางแผนการท่องเที่ยวให้คุ้มค่าและปลอดภัยยิ่งขึ้น

อย่างไรก็ตาม การใช้คะแนนความพึงพอใจเพียงอย่างเดียวในการประเมินคุณภาพสถานที่ท่องเที่ยวนั้นยังคงมีปัญหามากประการ เนื่องจากคะแนนที่ให้เป็นตัวเลขนับไม่สามารถบ่งบอกถึงข้อดีหรือข้อเสียที่แตกต่างกันในหลากหลายมุมมอง จึงไม่อาจสะท้อนถึงคุณภาพที่แท้จริงได้อย่างครบถ้วน [2] อีกทั้งความคาดหวังของนักท่องเที่ยวแต่ละคนนั้นไม่มีขีดจำกัด และยากที่จะเปรียบเทียบกันได้ด้วยหลักเกณฑ์ที่เป็นมาตรฐานเดียวกัน เพื่อแก้ไขปัญหาดังกล่าว จำเป็นต้องมีการวิเคราะห์ความรู้สึกจากบทวิจารณ์ (review) ที่อยู่ในรูปแบบภาษาธรรมชาติ เพื่อจำแนกออกเป็นความคิดเห็นเชิงบวก เชิงลบ และเป็นกลาง [3] หากแต่การวิเคราะห์ความรู้สึกจากข้อความเพียงอย่างเดียวอาจมีอคติจากผู้วิจารณ์ที่มีประสบการณ์สูงต่อทั้งเชิงบวกและเชิงลบ [4] ซึ่งทำให้เกิดความเอนเอียงและไม่ได้สะท้อนภาพรวมของสถานที่ท่องเที่ยวอย่างตรงไปตรงมา ส่งผลให้ข้อมูลที่เผยแพร่ขาดความน่าเชื่อถือและไม่เป็นกลาง ตลอดจนปริมาณของบทวิจารณ์จำนวนมากในปัจจุบัน ทำให้เกิดภาวะข้อมูลท่วมท้น (information overload) จนไม่สามารถแยกข้อมูลที่สำคัญได้ เนื่องจากไม่มีเครื่องมือหรือเทคนิคในการคัดกรองข้อมูลที่สำคัญออกมา ข้อมูลจำนวนมากทั้งที่เกี่ยวข้องและไม่เกี่ยวข้อง หรือข้อมูลที่ไม่จำเป็นต่าง ๆ อาจทำให้ผู้ใช้รู้สึกสับสนหรือ ไม่สามารถตัดสินใจได้อย่างเหมาะสม อีกทั้งสิ้นเปลืองเวลาและทรัพยากรในการแยกข้อมูลที่สำคัญ

หนึ่งในสถานที่ท่องเที่ยวที่ได้รับความนิยมมากที่สุดแห่งหนึ่งของประเทศไทยคืออุทยานแห่งชาติเขาใหญ่ ซึ่งเป็นพื้นที่มรดกโลกทางธรรมชาติ (World Heritage Site) และอุทยานมรดกแห่งอาเซียน (ASEAN Heritage Park) ครอบคลุมพื้นที่ 4 จังหวัด ประกอบด้วยนครราชสีมา ปราจีนบุรี นครนายก และสระบุรี ด้วยพื้นที่เกือบ 2,206 ตารางกิโลเมตร และเปิดต้อนรับนักท่องเที่ยวตลอดทั้งปี ในปี พ.ศ. 2565 มีจำนวนนักท่องเที่ยว

ถึง 1,428,765 คน แบ่งเป็นนักท่องเที่ยวชาวไทย 1,392,078 คน และชาวต่างชาติ 36,687 คน [5] ในปี พ.ศ. 2566 จำนวนของนักท่องเที่ยวเพิ่มเป็น 1,486,532 คน [6] โดยมีจำนวนเพิ่มขึ้นทั้งนักท่องเที่ยวชาวไทยและชาวต่างชาติ และมีแนวโน้มที่จะเพิ่มขึ้นเรื่อย ๆ ในอนาคต ด้วยธรรมชาติที่อุดมสมบูรณ์และความหลากหลายทางชีวภาพ จึงทำให้นักท่องเที่ยวจำนวนมากเดินทางมาเยี่ยมชมและแบ่งปันประสบการณ์ผ่านการเขียนบทวิจารณ์สาธารณะบนแพลตฟอร์มออนไลน์จำนวนมาก ทำให้ยากต่อการนำมาประมวลผลเพื่อใช้ประโยชน์ อีกทั้งมุมมองที่แตกต่างกันของนักท่องเที่ยวชาวไทยและชาวต่างชาติที่อาจมีความคาดหวังต่อประสบการณ์ที่แตกต่างกันของแหล่งท่องเที่ยว เช่นเดียวกับมุมมองของนักท่องเที่ยวชาวไทยอาจส่งผลทั้งทางตรงและทางอ้อมต่อความพึงพอใจในแหล่งท่องเที่ยวในประเทศอย่างมีนัยสำคัญ การนำเสนอข้อมูลจากทั้งสองกลุ่มนี้ในรูปแบบเดียวกันอาจทำให้การประเมินคุณภาพของสถานที่ท่องเที่ยวไม่ครอบคลุมทุกมิติ ดังนั้นการนำเสนอข้อมูลที่หลากหลายจากบทวิจารณ์ของนักท่องเที่ยวต่างชาติจะช่วยให้ผู้ที่เกี่ยวข้องเข้าใจความต้องการในมุมมองที่ครอบคลุมและนำไปสู่การพัฒนาคุณภาพให้ตรงตามความคาดหวังของนักท่องเที่ยวได้อย่างยั่งยืน

จากที่กล่าวมา คณะผู้วิจัยจึงได้เห็นถึงความสำคัญในการพัฒนากระบวนการวิเคราะห์ข้อมูลจากบทวิจารณ์ด้วยการประมวลผลภาษาธรรมชาติ และเครื่องการวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรม เพื่อนำเสนอข้อมูลที่ชัดเจนและตรงประเด็น นอกจากนั้นยังสะท้อนภาพรวมในแต่ละมิติของแหล่งท่องเที่ยวได้อย่างสมบูรณ์และเป็นธรรม อันจะเป็นหนึ่งในข้อมูลพื้นฐานในการตัดสินใจของนักท่องเที่ยว

2) ทบทวนวรรณกรรม

2.1) การประมวลผลภาษาธรรมชาติ (Natural Language Processing: NLP)

การประมวลผลภาษาธรรมชาติ เป็นการผสมผสานหลายศาสตร์เพื่อช่วยให้คอมพิวเตอร์สามารถเข้าใจและวิเคราะห์ข้อมูลในรูปแบบภาษาของมนุษย์ได้อย่างมีประสิทธิภาพ [7] ด้วยการสร้างระบบที่สามารถตีความข้อมูลเชิงภาษา ไม่ว่าจะเป็นข้อความหรือคำพูด และตอบสนองต่อมนุษย์ได้อย่างเหมาะสม

ปัจจุบันมีการประยุกต์ใช้เทคนิคการเรียนรู้เชิงลึก (deep learning) เพื่อเพิ่มประสิทธิภาพและความแม่นยำในการประมวลผลภาษาธรรมชาติ เช่น การแปลภาษาอัตโนมัติ (machine translation)

การสรุปความ (text summarization) เพื่อช่วยให้ผู้อ่านสามารถรับข้อมูลสำคัญได้อย่างรวดเร็ว และการรู้จำเสียงพูด (speech recognition) ซึ่งถูกนำไปใช้ในระบบผู้ช่วยอัจฉริยะ รวมถึงการวิเคราะห์ความรู้สึกที่มักถูกใช้ประเมินความคิดเห็นของผู้ใช้บนสื่อสังคมออนไลน์

2.2) การวิเคราะห์ข้อความ (Text Analysis)

การวิเคราะห์ข้อความ คือการประมวลผลเพื่อแปลความหมายจากข้อความด้วยเทคนิคและวิธีการทางคอมพิวเตอร์ [8] เพื่อนำไปใช้ในการสกัดข้อมูล การจัดหมวดหมู่ การสรุปเนื้อหา หรือการตรวจหาความสัมพันธ์ในข้อมูล แบ่งเป็นการวิเคราะห์ข้อความเชิงโครงสร้าง (structured text analysis) จากข้อมูลที่มีรูปแบบอย่างชัดเจน เช่น ข้อความในฐานข้อมูล หรือในเอกสารที่มีการกำหนดรูปแบบการนำเสนอ และการวิเคราะห์ข้อความไร้โครงสร้าง (unstructured text analysis) เช่น บทวิจารณ์ บทความ ข่าวสาร หรือความคิดเห็นของผู้ใช้บนสื่อสังคมออนไลน์ ซึ่งนำเสนออย่างอิสระ และไม่มีมาตรฐาน เนื่องจากอยู่ในรูปแบบของภาษาธรรมชาติ

2.3) การวิเคราะห์ความรู้สึก (Sentiment Analysis)

การวิเคราะห์ความรู้สึก เป็นกระบวนการที่เกี่ยวข้องกับการตีความและประเมินความคิดเห็นหรืออารมณ์จากข้อความในรูปแบบภาษาธรรมชาติ [3] เพื่อแยกแยะและประเมินความรู้สึกที่แสดงออกในข้อความนั้น ๆ ว่าเป็นเชิงบวก (positive) เชิงลบ (negative) หรือเป็นกลาง (neutral) จากการพิจารณาจากลักษณะของคำศัพท์ที่ใช้ในข้อความ โครงสร้างประโยค และการใช้สำนวนเฉพาะที่สะท้อนถึงอารมณ์ที่มีต่อเหตุการณ์หรือประเด็นต่าง ๆ ที่กล่าวถึง

งานวิจัยในปัจจุบันที่มุ่งวิเคราะห์ความรู้สึกจากข้อความบทวิจารณ์ ข้อความจากสื่อสังคมออนไลน์ หรือบทความข่าว สามารถจำแนกได้ 3 แนวทาง ประกอบด้วย 1) การวิเคราะห์ขั้วของความรู้สึก (polarity-based sentiment analysis) โดยจำแนกความรู้สึกออกเป็นเชิงบวก เชิงลบ และเป็นกลาง ซึ่งเป็นแนวทางที่ใช้กันอย่างแพร่หลายในปัจจุบัน เนื่องจากจุดเด่นที่ความเรียบง่ายและใช้งานกับข้อมูลขนาดใหญ่ได้อย่างมีประสิทธิภาพ เนื่องจากสามารถใช้พจนานุกรมอารมณ์สำเร็จรูปในการวิเคราะห์ข้อความได้เลยโดยไม่ต้องพัฒนาโมเดลที่ซับซ้อน นอกจากนี้ยังง่ายต่อการทำความเข้าใจและประยุกต์ใช้ได้อย่างหลากหลาย แต่ยังคงมีข้อจำกัดคือความแม่นยำ และอาจเกิดข้อผิดพลาดเมื่อต้อง

วิเคราะห์ข้อความที่ซับซ้อน เช่น การใช้ถ้อยคำเสียดสี (sarcasm) หรือการใช้คำในเชิงลบที่มีความหมายเชิงบวกในประโยค 2) การวิเคราะห์ระดับความเข้มของความรู้สึก (intensity-based sentiment analysis) ซึ่งเป็นการวิเคราะห์ที่มุ่งเน้นการวัดระดับของความรู้สึกในระดับที่ละเอียดมากขึ้น เช่น ความรู้สึกในระดับบวกมาก บวก เป็นกลาง ลบ และลบมาก แทนที่จะจำแนกเป็นเพียงเชิงบวกหรือเชิงลบเพียงอย่างเดียว ซึ่งจะสะท้อนถึงความรู้สึกที่ซ่อนอยู่ในข้อความได้มากขึ้น เหมาะสำหรับการวิเคราะห์ข้อมูลที่ต้องการติดตามความคิดเห็นในแต่ละช่วงเวลา หรือวิเคราะห์แนวโน้มของข้อความ อย่างไรก็ตาม การใช้งานจำเป็นต้องใช้ชุดข้อมูลที่มีการระบุระดับของความรู้สึกอย่างละเอียด นอกจากนั้นความแม่นยำของโมเดลขึ้นอยู่กับบริบทของข้อความ ซึ่งอาจทำให้เกิดข้อผิดพลาดในกรณีที่มีระดับความรู้สึกไม่ครอบคลุมทุกสถานการณ์ และ 3) การวิเคราะห์ประเภทของอารมณ์ (emotion-based sentiment analysis) โดยมุ่งเน้นการจำแนกความรู้สึกตามอารมณ์พื้นฐานของมนุษย์ในระดับที่เฉพาะเจาะจง เช่น สุข (joy) โศกเศร้า (sadness) โกรธ (anger) กลัว (fear) ประหลาดใจ (surprise) และขยะแขยง (disgust) เป็นต้น เนื่องจากสามารถระบุอารมณ์ที่เฉพาะเจาะจงได้มากกว่า จึงเหมาะสำหรับการวิเคราะห์เชิงจิตวิทยา หรือการประเมินภาวะสุขภาพจิตผ่านความคิดเห็นในสื่อสังคมออนไลน์ อย่างไรก็ตาม ข้อเสียของแนวทางนี้คือความซับซ้อนในการจัดเตรียมชุดข้อมูลที่มีป้ายกำกับอารมณ์ (emotion labeled) อย่างละเอียด และบางอารมณ์อาจมีความคล้ายคลึงกันจนยากต่อการแยกแยะ นอกจากนั้นต้องสร้างโมเดลขนาดใหญ่และต้องการทรัพยากรการประมวลผลสูง จึงมีข้อจำกัดในการใช้งานกับข้อมูลขนาดใหญ่ หรืออุปกรณ์ที่มีข้อจำกัดด้านทรัพยากร

กระบวนการวิเคราะห์ข้อความความรู้สึกเริ่มจากการเก็บรวบรวมข้อมูลจากแหล่งต่าง ๆ จากนั้นจึงเข้าสู่ขั้นตอนการเตรียมข้อมูลก่อนการประมวลผล (preprocessing) ซึ่งรวมถึงการลบอักขระพิเศษ การแปลงตัวอักษรเป็นตัวพิมพ์เล็ก การตัดคำ (tokenization) การปรับคำให้อยู่ในรูปแบบพื้นฐาน (lemmatization or stemming) และการลบคำหยุดหรือคำที่ไม่มีความหมายสำคัญ (stopword removal) ก่อนจะใช้การวิเคราะห์ความรู้สึกด้วยวิธีการต่าง ๆ ได้แก่ 1) การใช้พจนานุกรมอารมณ์ (lexicon-based approach) ที่มีการจัดอันดับความสำคัญของคำต่าง ๆ ไว้ในคลังคำศัพท์ [3] พร้อมด้วย ค่าความเข้มข้นของความรู้สึกหรืออารมณ์ (sentiment lexicon) เช่น เครื่องมือวิเคราะห์ความรู้สึกที่นิยมใช้กันมากใน

ปัจจุบัน เช่น TextBlob, Flair หรือ VADER ในการจับคู่คำศัพท์ และให้คะแนนความรู้สึก 2) การเรียนรู้ของเครื่อง (machine learning-based approach) [3] โดยอาศัยการสกัดคุณลักษณะ (feature extraction) จากการแปลงข้อความเป็นเวกเตอร์ เช่น TF-IDF ที่คำนึงถึงความสำคัญของคำในเอกสาร โดยคำที่มีความถี่สูงแต่ไม่สำคัญจะได้รับค่าน้ำหนักต่ำ รวมถึงการฝังคำ เช่น Word2Vec หรือ GloVe ที่แปลงคำให้เป็นเวกเตอร์ที่มีความหมาย ซึ่งสามารถจับความสัมพันธ์ระหว่างคำที่มีความหมายคล้ายกันได้ ตลอดจนเทคนิค one-hot encoding ซึ่งแปลงคำให้เป็นเวกเตอร์ที่มีมิติเท่ากับจำนวนคำทั้งหมดในพจนานุกรม โดยที่ทุกตำแหน่งในเวกเตอร์จะเป็นศูนย์ ยกเว้นตำแหน่งที่ตรงกับคำที่กำลังแปลงจะมีค่าเป็น 1 ก่อนจะใช้อัลกอริทึมการเรียนรู้ของเครื่องทั้งแบบมีผู้สอนและไม่มีผู้สอน เช่น ซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) [9] นาอิวเบย์ (Naive Bayes) หรือการจัดกลุ่มแบบเคมีน (k-mean clustering) [10] และ 3) การเรียนรู้เชิงลึก (deep learning-based approach) [11] จากการนำโครงข่ายประสาทเทียมมาประยุกต์ใช้ เช่น โครงข่ายประสาทเทียมแบบวนซ้ำ (Recurrent Neural Network: RNN) หน่วยความจำระยะสั้นแบบยาว (Long Short-Term Memory: LSTM) และโมเดลการเรียนรู้เชิงลึกด้านภาษาอย่าง transformer-based models [12] เช่น BERT หรือ DistilBERT ที่แยกแยะบริบทของคำศัพท์และความหมายแฝงในข้อความได้ดียิ่งขึ้น ทำให้สามารถระบุประเด็นที่สำคัญได้อย่างมีประสิทธิภาพ

เมื่อพิจารณาแต่ละแนวทางจะพบว่ามีข้อดีและข้อเสียแตกต่างกันไป ดังนั้นการเลือกใช้แต่ละวิธีการจึงขึ้นอยู่กับลักษณะของข้อมูล ความแม่นยำที่ต้องการ และทรัพยากรการประมวลผล โดยโมเดลการเรียนรู้ของเครื่องนั้นมีข้อดีในด้านความเร็วและความง่ายในการฝึกฝนโมเดล แต่ในขณะเดียวกันก็มีข้อจำกัดเมื่อข้อมูลที่ใช้นั้นมีความซับซ้อนหรือมีลำดับคำที่สำคัญชัดเจน ในทางกลับกันการใช้โมเดลการเรียนรู้เชิงลึกที่สามารถจัดการกับข้อมูลลำดับได้นั้น [13] แต่ยังมีข้อเสียในด้านความซับซ้อนของโมเดล อีกทั้งการฝึกฝนโมเดลที่ใช้เวลานาน และต้องการทรัพยากรการคำนวณสูง จึงอาจทำให้เกิดข้อจำกัดในกรณีที่ต้องการ การประมวลผลแบบเรียลไทม์ เป็นที่มาของการเลือกใช้เครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมบนเงื่อนไขที่ต้องการเน้นความเร็วในการประมวลผลและใช้ทรัพยากรคำนวณต่ำ แต่ยังสามารถวิเคราะห์ความรู้สึกจากข้อความได้อย่างแม่นยำ

2.4) เครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรม (Lexicon-Based Sentiment Analysis Tool)

เครื่องมือวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรมส่วนใหญ่อ้างอิงจากฐานข้อมูลที่บรรจุคำศัพท์พร้อมค่าคะแนนความรู้สึก ซึ่งมีจุดเด่นคือง่ายต่อการใช้งานโดยไม่ต้องฝึกฝนโมเดลเหมือนวิธีการเรียนรู้ของเครื่อง และสามารถวิเคราะห์ข้อมูลจากข้อความในลักษณะของคำและประโยค โดยไม่จำเป็นต้องแปลงเป็นเวกเตอร์ตัวเลขก่อน โดยเครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมที่ได้รับความนิยมในปัจจุบัน ได้แก่

2.4.1) TextBlob [14] เป็นไลบรารีที่พัฒนาขึ้นเพื่อรองรับการวิเคราะห์ความรู้สึกจากข้อความโดยใช้คลังคำศัพท์ที่กำหนดค่าคะแนนความรู้สึกที่กำหนดไว้ล่วงหน้า ทำให้สามารถระบุได้ว่าข้อความมีแนวโน้มความรู้สึกเชิงบวก เชิงลบ หรือเป็นกลางจากการใช้ pattern library ซึ่งเป็นเครื่องมือที่ใช้สำหรับการทำเหมืองข้อความและวิเคราะห์โครงสร้างของภาษา กระบวนการวิเคราะห์ความรู้สึกมีพื้นฐานบน 2 องค์ประกอบหลัก ได้แก่ ค่าความเป็นบวกหรือลบ (polarity) ถูกกำหนดอยู่ในช่วงระหว่าง -1 ถึง +1 โดยค่าที่เข้าใกล้ +1 แสดงถึงข้อความที่มีความรู้สึกเชิงบวกสูง ในทางกลับกันค่าที่เข้าใกล้ -1 หมายถึงข้อความมีความรู้สึกเชิงลบสูง ร่วมกับค่าอัตวิสัย (subjectivity) ที่แสดงถึงความคิดเห็นส่วนตัวหรือข้อเท็จจริงที่ถูกกำหนดให้อยู่ในช่วง 0 ถึง 1 โดยค่าที่เข้าใกล้ 1 แสดงถึงข้อความที่เป็นความคิดเห็นส่วนตัวมากกว่าข้อเท็จจริง ในขณะที่ค่าที่ใกล้ 0 บ่งชี้ว่าข้อความนั้นมีลักษณะเป็นกลางและมีความเป็นข้อเท็จจริงมากกว่า

ความง่ายในการใช้งานและไม่จำเป็นต้องฝึกโมเดลเพิ่มเติมก่อนนำไปใช้วิเคราะห์ความรู้สึกจากข้อความถือเป็นจุดเด่นสำคัญของ TextBlob นอกจากนี้ยังรองรับการประมวลผลภาษาธรรมชาติในด้านอื่น ๆ เช่น การจำแนกข้อความ การวิเคราะห์ไวยากรณ์ (part-of-speech tagging) การแยกคำนาม (noun phrase extraction) และการแปลภาษา โดยเฉพาะการทำงานกับความคิดเห็นที่มีการใช้ภาษาทางการในรูปแบบภาษาอังกฤษ อย่างไรก็ตาม ยังมีข้อจำกัดในการวิเคราะห์ข้อความที่มีความหมายซับซ้อน เนื่องจากอาศัยกฎทางภาษาศาสตร์และคะแนนเชิงอารมณ์ที่ถูกกำหนดล่วงหน้าเพียงอย่างเดียว เช่นเดียวกับประสิทธิภาพอาจลดลงเมื่อใช้งานกับข้อความที่มีการใช้ภาษาเชิงประชดประชัน หรือข้อความที่มีบริบทซับซ้อนที่ต้องอาศัยความเข้าใจทางสังคมและวัฒนธรรม

2.4.2) *Flair* [15] เป็นหนึ่งในเครื่องมือการวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมด้วยเทคนิคการฝังคำตามบริบท (contextualized word embeddings) และใช้หน่วยความจำระยะสั้นแบบยาวสองทิศทาง (bidirectional long short-term memory: BiLSTM) ที่ผ่านการฝึกฝนล่วงหน้าด้วยข้อมูลจำนวนมาก ซึ่งช่วยให้สามารถจับความสัมพันธ์ระหว่างคำในประโยค และทำความเข้าใจความหมายของข้อความจากทั้งสองทิศทางได้อย่างมีประสิทธิภาพ โดยในกระบวนการนี้โมเดลจะการเรียนรู้จากบริบทของข้อความในประโยคก่อนหน้าและหลังเพื่อระบุความหมายของคำเป้าหมาย ทำให้สามารถประเมินอารมณ์จากข้อความได้แม่นยำยิ่งขึ้น จึงเหมาะกับการวิเคราะห์ข้อความที่มีโครงสร้างซับซ้อนหรืองานที่ต้องการตีความหมายในระดับสูง เช่น ข้อความที่ใช้ภาษาเชิงประชดประชัน เป็นต้น

จุดเด่นสำคัญของ Flair คือเป็นโมเดลที่ฝึกฝนล่วงหน้าจากหลายภาษา จึงสามารถทำงานได้อย่างหลากหลายไม่จำกัดเฉพาะภาษาอังกฤษ และสามารถนำไปใช้งานได้ทันทีโดยไม่ต้องฝึกโมเดลใหม่ จึงมักถูกนำไปใช้ในการระบุองค์ประกอบในประโยค (named entity recognition) การระบุประเภทข้อความ (text classification) การตอบคำถามอัตโนมัติ (question answering) และการสรุปความ (summarization) แต่ก็มีข้อจำกัดคือการใช้ทรัพยากรและเวลาในการประมวลผลที่มากกว่า TextBlob และ VADER จึงไม่เหมาะสมกับการทำงานบนอุปกรณ์ที่มีข้อจำกัดด้านทรัพยากร

2.4.3) *Valence Aware Dictionary and Sentiment Reasoner: VADER* [16] เป็นไลบรารีที่ถูกออกแบบสำหรับวิเคราะห์ความรู้สึกในข้อความภาษาธรรมชาติโดยเฉพาะ จากการประยุกต์ใช้พจนานุกรมที่มีการจัดอันดับความสำคัญของคำต่าง ๆ มากกว่า 7,500 คำ ซึ่งแต่ละคำถูกกำหนดคะแนนความรู้สึก (sentiment score) ที่สะท้อนถึงความเข้มข้นของอารมณ์หรือความรู้สึกไว้ล่วงหน้า ดังตารางที่ 1

ตารางที่ 1 : คำศัพท์และตัวอย่างค่าคะแนนความรู้สึกในพจนานุกรม

word	sentiment score
happy	+2.9
sad	-2.1
excellent	+3.2
horrible	-3.1
not bad	+1.0

จากตารางที่ 1 คำศัพท์ในพจนานุกรมที่มีกำหนดตัวอย่างคะแนนความรู้สึกจะถูกนำมาใช้วิเคราะห์ข้อความร่วมกับการกำหนดฐานกฎ (rule-based) ที่กำหนดไว้ล่วงหน้าจะช่วยเพิ่มความแม่นยำในการพิจารณาลักษณะทางภาษา เพื่อตัดสินใจว่าความรู้สึกในข้อความนั้นเป็นเชิงบวก เชิงลบ หรือเป็นกลาง แทนที่จะให้ค่าความสัมพันธ์ของความรู้สึกแต่ละคำเพียงอย่างเดียวแบบเครื่องมือวิเคราะห์ความรู้สึกจากข้อความทั่วไป

หนึ่งในการพิจารณาลักษณะทางภาษาคือการกำหนดลักษณะเฉพาะที่ส่งผลต่อค่าคะแนนอารมณ์ เช่น หากใช้ตัวอักษรพิมพ์ใหญ่ทั้งหมด เช่น “AWESOME” จะกำหนดค่าความรู้สึกรุนแรงกว่า “awesome” ที่ใช้ตัวอักษรพิมพ์เล็กทั้งหมด นอกจากนี้เครื่องหมายอัศเจรีย์ (!) ยังช่วยเพิ่มระดับความเข้มของอารมณ์ เช่น “great!!” ให้ความรู้สึกเชิงบวกมากกว่าคำว่า “great” แบบปกติ ตลอดจนสามารถพิจารณาคำขยาย (modifiers) ที่มีผลต่อระดับความเข้มของความรู้สึกหรือคำวิเศษณ์ เช่น “very extremely” และ “slightly” ที่สามารถเปลี่ยนระดับความรุนแรงของคำได้ เช่นเดียวกับ “very good” จะได้รับคะแนนเชิงบวกสูงกว่าคำว่า “good” ในทางตรงกันข้าม “not good” จะถูกลดระดับความรู้สึกเชิงบวกของคำว่า “good” ทำให้ความหมายปรับลดไปเป็นกลางหรือเชิงลบมากขึ้น เช่นเดียวกับ “not bad” จะถูกลดระดับความรู้สึกเชิงลบจากคำว่า “bad” และปรับความรู้สึกให้เข้าใกล้กับคำว่า “good” ยิ่งขึ้น

แม้ว่าจะยังมีข้อจำกัดบางประการ แต่ VADER ยังสามารถตรวจจับคำประชดประชันได้ในระดับที่น่าพอใจ โดยอาศัยบริบทของข้อความ เช่น วลี “yeah, right” ที่มักใช้ในเชิงประชดประชัน ซึ่งสามารถนำมาพิจารณาเพื่อลดโอกาสในการแปลความหมายผิดพลาด ตลอดจนความสามารถในการจัดการกับการใช้ภาษาที่ไม่เป็นทางการ เช่น ข้อความที่มีการย่อคำ การใช้ตัวอักษรแทนการสื่ออารมณ์ (emoticon) และการใช้คำเสียดสีหรือแสดงความขัดแย้งในตัวเองได้ นอกจากนี้ยังมีความเร็วในการประมวลผลและไม่ต้องการทรัพยากรคำนวณมากนัก จึงเหมาะสำหรับการใช้งานที่ต้องการการวิเคราะห์แบบเรียลไทม์ หรือในระบบที่มีข้อจำกัดด้านทรัพยากร ตลอดจนคงทนต่อการเปลี่ยนแปลงของข้อมูลได้โดยไม่มีข้อจำกัดเรื่องมาตรฐาน จึงเหมาะสำหรับวิเคราะห์ความรู้สึกในข้อความที่ไม่เป็นทางการ จากที่กล่าวมาจึงเป็นเหตุผลหลักที่ทำให้ VADER ถูกนำมาใช้ในการประมวลผลข้อความที่เป็นภาษาธรรมชาติใน สื่อสังคมออนไลน์อย่างแพร่หลายในปัจจุบัน

2.5) บทวิจารณ์ออนไลน์ (Online Review)

บทวิจารณ์ออนไลน์ คือความคิดเห็นที่สะท้อนมุมมองหรือประสบการณ์ที่เคยได้รับผ่านแพลตฟอร์มดิจิทัล โดยบทวิจารณ์ที่เหล่านี้นักเป็นประสบการณ์ตรงจากผู้ที่เคยเดินทางท่องเที่ยว ณ สถานที่แห่งนั้นในช่วงเวลาที่ผ่านไป ซึ่งสะท้อนถึงคุณภาพของแหล่งท่องเที่ยว ที่พัก ร้านอาหาร และบริการต่าง ๆ ที่เกี่ยวข้อง เผยแพร่บนแพลตฟอร์มออนไลน์ เช่น Google reviews, Agoda.com, TripAdvisor.com, Trip.com และ สื่อสังคมออนไลน์ต่าง ๆ ซึ่งมีบทบาทในการให้ข้อมูลและกำหนดพฤติกรรมของนักท่องเที่ยว เช่น การเปรียบเทียบที่พักจากบทวิจารณ์ออนไลน์ อย่างไรก็ดี แม้ว่าบทวิจารณ์ออนไลน์จะมีข้อดีหลายประการ แต่ก็ยังมีข้อจำกัดที่ต้องพิจารณา เช่น ความรู้สึกส่วนตัวที่มีอิทธิพลต่อการเขียนบทวิจารณ์ ดังนั้นการวิเคราะห์ข้อมูลอย่างโปร่งใสและสร้างสรรค์จะช่วยยกระดับคุณภาพของการท่องเที่ยวได้อย่างยั่งยืน

2.6) องค์ประกอบแหล่งท่องเที่ยว

การพัฒนาแหล่งท่องเที่ยวสามารถแบ่งออกได้เป็นหลายองค์ประกอบที่ช่วยส่งเสริมให้แหล่งท่องเที่ยวนั้นมีความน่าสนใจและกระตุ้นให้เกิดการเดินทางมายังสถานที่นั้น ๆ ได้ โดยอุทยานแห่งชาติเขาใหญ่เป็นแหล่งท่องเที่ยวที่มีชื่อเสียงระดับโลก มีองค์ประกอบที่สอดคล้องกับกรอบการวิเคราะห์จุดหมายปลายทางการท่องเที่ยวแบบ 6A (6A's framework for the analysis of tourism destinations) ของ Buhalis [17] ซึ่งแบ่งออกเป็น 6 ด้าน ได้แก่ 1) สิ่งดึงดูดใจ (attractions) อุทยานแห่งชาติเขาใหญ่มีสิ่งดึงดูดใจที่โดดเด่นในหลากหลายมิติ ได้แก่ สถานที่ท่องเที่ยวเชิงธรรมชาติ เช่น น้ำตกเหวสุวัต น้ำตกเหวนรก จุดชมวิวเขาเขียว และผืนป่าอันอุดมสมบูรณ์ที่เต็มไปด้วยพืชพรรณและสัตว์ป่าหลากหลายชนิด สถานที่ท่องเที่ยวที่มนุษย์สร้าง เช่น เส้นทางศึกษาธรรมชาติและอ่างเก็บน้ำสายสร สถานที่ท่องเที่ยวเชิงศิลปวัฒนธรรม เช่น ประวัติศาสตร์และวัฒนธรรมของชุมชนรอบอุทยานที่สะท้อนถึงวิถีชีวิตและความสัมพันธ์ระหว่างมนุษย์กับธรรมชาติ และสถานที่ท่องเที่ยวเชิงชุมชนสัมพันธ์ เช่น การเยี่ยมชมตลาดพื้นเมือง หรือการเข้าร่วมกิจกรรมเชิงอนุรักษ์ธรรมชาติในพื้นที่ 2) การเข้าถึงสถานที่ท่องเที่ยว (accessibility) การเดินทางมายังอุทยานแห่งชาติเขาใหญ่ผ่านถนนสายหลักที่เชื่อมต่อกันได้หลายเส้นทาง โดยเฉพาะการเดินทางด้วยรถยนต์ส่วนบุคคล หรือรถโดยสารสาธารณะ นอกจากนี้การลงทะเบียน

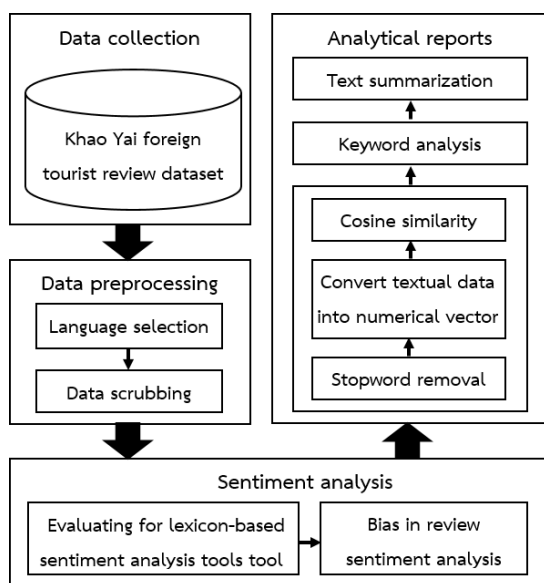
เข้าอุทยานล่วงหน้าผ่าน แอปพลิเคชัน และการจองที่พักผ่านเว็บไซต์และช่องทางออนไลน์ต่าง ๆ ก็มีส่วนช่วยให้นักท่องเที่ยวสามารถวางแผนการเดินทางได้อย่างมีประสิทธิภาพ 3) สิ่งอำนวยความสะดวก (amenities) เพื่อรองรับนักท่องเที่ยว เช่น ที่พักแบบบ้านพักและลานกางเต็นท์ ร้านอาหารภายในพื้นที่ และศูนย์บริการนักท่องเที่ยว นอกจากนี้ยังมีบริการด้านสุขอนามัย เช่น ห้องน้ำสาธารณะ และบริการสาธารณสุขประเภทต่าง ๆ 4) การให้บริการแบบเป็นชุด (available packages) ที่ตอบสนองความต้องการของนักท่องเที่ยว เช่น แพคเกจเดินป่า แพคเกจชมสัตว์ป่าในช่วงกลางคืน หรือแพคเกจเยี่ยมชมน้ำตกภายในพื้นที่ อีกทั้งยังมีการจัดกิจกรรมพิเศษในช่วงฤดูกาลต่าง ๆ เช่น การชมนกในฤดูหนาว หรือการเดินป่าศึกษาธรรมชาติในช่วงฤดูฝน เป็นต้น 5) กิจกรรมการท่องเที่ยว (activities) ที่หลากหลายและตอบสนองความสนใจของนักท่องเที่ยวกลุ่มต่าง ๆ ที่เดินทางมาในรูปแบบเดี่ยว แบบคู่รัก แบบครอบครัว ตลอดจนแบบหมู่คณะ ไม่ว่าจะเป็นการเดินป่าศึกษาธรรมชาติ การปั่นจักรยานในเส้นทางธรรมชาติ การส่องสัตว์ป่า และการถ่ายภาพทิวทัศน์ กิจกรรมเหล่านี้ช่วยเสริมสร้างประสบการณ์ที่น่าจดจำให้แก่นักท่องเที่ยว และ 6) การให้บริการของแหล่งท่องเที่ยว (ancillary services) อุทยานแห่งชาติ เขาใหญ่ยังมีการจัดเตรียมบริการเสริมที่สำคัญ เช่น หน่วยบริการทางการแพทย์เบื้องต้น ธนาคาร บริการชำระเงินอิเล็กทรอนิกส์ บริการโทรศัพท์และอินเทอร์เน็ต รวมถึงจุดบริการข้อมูลที่ช่วยอำนวยความสะดวกแก่นักท่องเที่ยว จะเห็นได้ว่าองค์ประกอบทั้ง 6 ด้านนี้แสดงให้เห็นถึงศักยภาพของอุทยานแห่งชาติเขาใหญ่ในการตอบสนองความต้องการของนักท่องเที่ยวได้อย่างครบถ้วน

อย่างไรก็ดี เนื่องจากบริบทของอุทยานแห่งชาติเขาใหญ่เป็นแหล่งท่องเที่ยวทางธรรมชาติ ดังนั้นการท่องเที่ยวส่วนใหญ่จึงมุ่งเน้นการอนุรักษ์ธรรมชาติและสิ่งแวดล้อมเป็นหลัก งานวิจัยนี้จึงพิจารณารวมการให้บริการของแหล่งท่องเที่ยวมารวมกับองค์ประกอบด้านสิ่งอำนวยความสะดวก เนื่องจากเป็นบริการเสริมที่เกี่ยวข้องทั้งแบบเป็นทางการ และไม่เป็นทางการที่นักท่องเที่ยวควรได้รับ เช่นเดียวกับการให้บริการแบบเป็นชุดที่จะมีความทับซ้อนกับกิจกรรมการท่องเที่ยวในพื้นที่ คณะผู้วิจัยจึงรวมเข้ากับองค์ประกอบด้านกิจกรรม เนื่องจากเป็นบริบทที่มีความใกล้เคียงกัน เพื่อให้ได้ข้อมูลพื้นฐานที่ใช้เป็นแนวทางในการปรับปรุง และพัฒนาให้ตอบสนองต่อความต้องการของนักท่องเที่ยว ซึ่งจะนำไปสู่การ

พัฒนาอุตสาหกรรมการท่องเที่ยวอย่างยั่งยืนในจังหวัดนครราชสีมา และพื้นที่ใกล้เคียง

3) วิธีดำเนินการวิจัย

งานวิจัยนี้ประยุกต์ใช้การประมวลผลภาษาธรรมชาติร่วมกับเครื่องมือการวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรมเพื่อวิเคราะห์บทวิจารณ์ของนักท่องเที่ยวชาวต่างชาติต่ออุทยานแห่งชาติเขาใหญ่ แบ่งการดำเนินงานเป็น 4 ขั้นตอน ดังรูปที่ 1 กรอบการดำเนินงาน ประกอบด้วย การรวบรวมข้อมูล การประมวลผลข้อมูลเบื้องต้น การวิเคราะห์ความรู้สึก และการสร้างรายงานเชิงวิเคราะห์



รูปที่ 1 : กรอบการดำเนินงาน

3.1) การรวบรวมข้อมูล (Data Collection)

การรวบรวมบทวิจารณ์สาธารณะของนักท่องเที่ยวต่ออุทยานแห่งชาติเขาใหญ่ด้วยการสกัดและดึงข้อมูลจากเว็บไซต์ (web scraping) [18] ที่เปิดให้เขียนบทวิจารณ์สาธารณะ ตั้งแต่วันที่ 1 มกราคม พ.ศ. 2563 ถึง 31 ธันวาคม พ.ศ. 2567 เนื่องจากข้อมูลดังกล่าวสะท้อนถึงประสบการณ์ และมุมมองของนักท่องเที่ยวในช่วงเวลาปัจจุบัน

อย่างไรก็ตาม งานวิจัยนี้ให้ความสำคัญต่อข้อกำหนดทางกฎหมายและจริยธรรมในการดึงข้อมูลจากแพลตฟอร์มออนไลน์ โดยยึดมั่นในหลักการปฏิบัติตามนโยบายและข้อกำหนดที่ระบุโดยเว็บไซต์ที่ใช้เป็นแหล่งข้อมูล นอกจากนี้ยังตระหนักถึงความสำคัญของการปกป้องข้อมูลส่วนบุคคลของผู้ใช้ ซึ่งรวมถึงการ

หลีกเลี่ยงข้อมูลที่สามารถระบุตัวบุคคลได้โดยตรงที่อาจสร้างความเสี่ยงต่อความเป็นส่วนตัวของผู้วิจารณ์ รวมถึงภาพถ่ายหรือสื่ออื่น ๆ ที่มีลิขสิทธิ์ของผู้ใช้มาใช้เป็นส่วนหนึ่งของการวิเคราะห์เพื่อให้มั่นใจว่างานวิจัยดำเนินการภายใต้กรอบของความถูกต้องตามกฎหมายและจริยธรรมที่ไม่กระทบต่อความเป็นส่วนตัวหรือสิทธิของผู้ใช้งานแพลตฟอร์มออนไลน์ในทางใดทางหนึ่ง

3.2) การประมวลผลข้อมูลเบื้องต้น (Data Preprocessing)

หลังจากที่ได้ข้อมูลดิบจากการรวบรวมข้อมูลแล้ว ขั้นตอนต่อมาคือการประมวลผลข้อมูลเบื้องต้น เพื่อแปลงข้อมูลให้อยู่ในรูปแบบที่พร้อมสำหรับการวิเคราะห์ มีกระบวนการดังนี้

3.2.1) การคัดเลือกภาษา (Language Selection) เนื่องจากบทวิจารณ์สาธารณะมักประกอบด้วยหลายภาษา งานวิจัยนี้มุ่งวิเคราะห์บทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติ จึงคัดกรองเฉพาะบทวิจารณ์ที่เขียนด้วยภาษาอังกฤษที่เกี่ยวข้องกับการท่องเที่ยวในอุทยานเท่านั้น เพื่อให้ได้ข้อมูลที่มีความเฉพาะเจาะจงและสอดคล้องกับวัตถุประสงค์ของการศึกษา

3.2.2) การทำความสะอาดข้อมูล (Data Scrubbing) ก่อนการนำไปใช้งาน โดยตรวจสอบข้อมูลที่สูญหาย (missing values) ซึ่งอาจมีผลกระทบต่อความถูกต้องของการวิเคราะห์ นอกจากนี้ยังมีการตรวจสอบและลบข้อมูลที่ไม่เกี่ยวข้อง ซึ่งอาจรวมถึงข้อมูลที่ไม่มีความหมายหรือเป็นข้อมูลรบกวน เช่น สัญลักษณ์ หรืออักขระพิเศษที่ไม่เกี่ยวข้อง ตลอดจนการปรับเปลี่ยนรูปแบบของข้อมูลให้ข้อมูลมีความสม่ำเสมอและพร้อมสำหรับการประมวลผล

อย่างไรก็ดี เนื่องจากบทวิจารณ์อยู่ในรูปแบบของภาษาอังกฤษ ซึ่งมีการเว้นวรรคเพื่อระบุขอบเขตของคำอย่างชัดเจน จึงไม่จำเป็นต้องผ่านกระบวนการตัดคำ ซึ่งจะช่วยให้การประมวลผลเป็นไปอย่างรวดเร็วและมีประสิทธิภาพมากขึ้น

3.3) การวิเคราะห์ความรู้สึก (Sentiment Analysis)

การวิเคราะห์ความรู้สึกจากบทวิจารณ์ในรูปแบบภาษาอังกฤษจำนวน 12,035 รายการที่ผ่านการประมวลผลข้อมูลเบื้องต้น รายละเอียดดังตารางที่ 2

ตารางที่ 2 : จำนวนบทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติ

Rating	Quantity (items)
1	546
2	560
3	1,644

ตารางที่ 2 : จำนวนบทวิจารณ์จากนักท่องเที่ยวต่างชาติ (ต่อ)

Rating	Quantity (items)
4	2,814
5	6,471
รวม	12,035

จากตารางที่ 2 บทวิจารณ์ที่ผ่านการประมวลผลข้อมูลเบื้องต้นแล้วจะนำมาวิเคราะห์ความรู้สึกจากข้อความร่วมกับค่าคะแนนความพึงพอใจ 5 ระดับ เพื่อประเมินผลด้วยค่าความถูกต้องร่วมกับการวิเคราะห์ความเอนเอียงในบทวิจารณ์ มีรายละเอียดดังนี้

3.3.1) การประเมินผลเครื่องมือสำหรับวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรม (*Evaluating for Lexicon-Based Sentiment Analysis Tools*) ที่สามารถวิเคราะห์ข้อความในรูปแบบข้อความธรรมดาได้เลย โดยไม่จำเป็นต้องแปลงเป็นเวกเตอร์เชิงตัวเลขก่อน

งานวิจัยนี้ประเมินผลด้วยค่าความถูกต้อง (accuracy) เพื่อวัดประสิทธิภาพของเครื่องมือสำหรับวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมที่ได้แก่ TextBlob, Flair และ VADER ในการวิเคราะห์และจำแนกความรู้สึกจากข้อความ ดังสมการที่ 1

$$Accuracy = \frac{\text{Number of correct predictions}}{\text{Total number of test instances}} \quad (1)$$

จากสมการที่ 1 ค่าความถูกต้องจะคำนวณจากสัดส่วนของการจำแนกข้อมูลที่ทดสอบว่าเป็นเชิงบวก เชิงลบ หรือเป็นกลาง เมื่อเทียบกับผลจริงที่ได้จากการให้คะแนนความพึงพอใจ โดยกำหนดให้ค่าคะแนนความพึงพอใจที่เท่ากับ 5 และ 4 ถือว่าเป็นความรู้สึกเชิงบวก หากเท่ากับ 3 ถือว่าเป็นกลาง และค่าคะแนนความพึงพอใจที่เท่ากับ 2 และ 1 ถือว่าเป็นความรู้สึกเชิงลบ คิดเป็นจำนวนผลลัพธ์จริงของความรู้สึกเชิงบวกเท่ากับ 9,285 รายการ เชิงลบ เท่ากับ 1,106 รายการ และเป็นกลางเท่ากับ 1,644 รายการ ดังนั้นค่าความถูกต้องที่ได้จะแสดงให้เห็นถึงความสามารถในการจำแนกความรู้สึก และเป็นแนวทางในการนำไปประยุกต์ใช้งานในสถานการณ์จริงต่อไปในอนาคต

3.3.2) การวิเคราะห์ความเอนเอียงในบทวิจารณ์ (*Bias in Review Sentiment Analysis*) เนื่องจากในการวิเคราะห์ข้อความหนึ่งในแนวคิดที่ได้รับความนิยมอย่างมากคือการประเมินความรู้สึกจากบทวิจารณ์ออนไลน์ อย่างไรก็ตาม ปัญหาหนึ่งที่มักเกิดขึ้นคือความเอนเอียง (bias) ในบทวิจารณ์ ซึ่งทำให้ผลลัพธ์ของการวิเคราะห์มีความคลาดเคลื่อนจากความเป็นจริง

ความเอนเอียงในบทวิจารณ์อาจเกิดจากปัจจัยหลายประการ เช่น แรงจูงใจในการให้คะแนน โดยบางแพลตฟอร์มให้รางวัลหรือมีสิทธิพิเศษสำหรับบทวิจารณ์เชิงบวก ซึ่งอาจทำให้คะแนนความพึงพอใจมีแนวโน้มสูงกว่าความเป็นจริง ตลอดจนความเข้าใจผิดเกี่ยวกับระบบการให้คะแนน อารมณ์ชั่วขณะ หรือการใช้ภาษาที่มีลักษณะประชดประชัน งานวิจัยนี้จึงมุ่งวิเคราะห์ความเอนเอียงจากบทวิจารณ์ที่มีคะแนนความพึงพอใจเท่ากับ 5 แต่มีความรู้สึกจากข้อความเป็นเชิงลบ และในทางกลับกัน บทวิจารณ์ที่มีคะแนนความพึงพอใจเท่ากับ 1 แต่มีความรู้สึกจากข้อความเป็นเชิงบวกจากการใช้เครื่องมือวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรม โดยแต่ละค่าจะได้รับคะแนนความเข้มข้น (valence score) [19] ที่จะช่วยให้ระบุทิศทางของอารมณ์ในข้อความ หากมีคะแนนมากกว่า 0.05 จะถูกพิจารณาว่าเป็นเชิงบวก และหากคะแนนน้อยกว่า -0.05 จะถูกพิจารณาว่าเป็นเชิงลบ นอกจากนี้หากคะแนนอยู่ระหว่าง -0.05 ถึง 0.05 จะถูกพิจารณาว่าเป็นกลาง และคะแนนรวม (compound score) ที่พิจารณาความรู้สึกทั้งหมดของข้อความที่ในบทวิจารณ์ เนื่องจากในบทวิจารณ์เดียวกันอาจมีการผสมผสานของความคิดเห็นทั้งบวกและลบ คะแนนรวมจึงช่วยให้สามารถประเมินความรู้สึกโดยรวมของข้อความได้อย่างแม่นยำจากการรวมคะแนนของค่าต่าง ๆ ที่ปรากฏในบทวิจารณ์

คะแนนรวมที่มีเกณฑ์ในการตัดสินใจที่ชัดเจนสามารถนำไปหาความเอนเอียงระหว่างคะแนนความพึงพอใจเปรียบเทียบกับผลการวิเคราะห์ความรู้สึก เช่น หากคะแนนความพึงพอใจสูงแต่แสดงความรู้สึกเชิงลบในข้อความ หมายความว่านักท่องเที่ยวอาจมีความคาดหวังที่สูง หรืออาจมีข้อร้องเรียนบางประการที่ไม่สามารถสะท้อนในคะแนนได้ ซึ่งแสดงให้เห็นถึงปัญหาที่ต้องได้รับการแก้ไข ในทางตรงกันข้าม คะแนนความพึงพอใจต่ำแต่กลับแสดงความรู้สึกเชิงบวกในข้อความ อาจบ่งชี้ถึงความไม่ชัดเจนในระบบการประเมิน หรือการให้คะแนนที่อาจตีความไม่ตรงกัน จึงต้องตรวจสอบวิธีการให้คะแนนหรือนโยบายการรวบรวมข้อมูลเพิ่มเติม

3.4) การสร้างรายงานเชิงวิเคราะห์ (Analytical Reports)

รายงานเชิงวิเคราะห์เพื่อมุ่งเน้นการนำเสนอข้อมูลเชิงลึกในรูปแบบที่เข้าใจง่ายจากการเชื่อมโยงข้อมูลเพื่อวิเคราะห์ในหลากหลายมิติ ร่วมกับการสื่อสารผลลัพธ์ผ่านการสร้างภาพข้อมูล (data visualization) ประกอบด้วย

3.4.1) การลบคำหยุด (Stop Removal) ที่พบได้บ่อยครั้งในข้อความแต่ไม่มีนัยสำคัญ จึงสามารถตัดออกได้เลยโดยไม่กระทบต่อความหมายหลักของข้อความ เช่น “a”, “the”, “is” หรือ “and” เป็นต้น ซึ่งช่วยลดขนาดของข้อมูลและทำให้สามารถมุ่งเน้นการวิเคราะห์ไปยังคำที่มีความหมายสำคัญได้มากขึ้น จากนั้นจะลบเครื่องหมายต่าง ๆ ที่ไม่เกี่ยวข้อง ก่อนจะแปลงข้อความทั้งหมดเป็นตัวอักษรพิมพ์เล็กในรูปแบบพื้นฐาน เช่น “watching” เปลี่ยนเป็น “watch” ขั้นตอนนี้ช่วยลดความหลากหลายของคำที่มีรูปแบบแตกต่างกันแต่มีความหมายเดียวกัน ให้ข้อมูลมีความสม่ำเสมอและพร้อมสำหรับการประมวลผล

3.4.2) การแปลงข้อความเป็นเวกเตอร์เชิงตัวเลขจากการพิจารณาความสำคัญของคำในแต่ละข้อความผ่านการคำนวณ TF-IDF (Term Frequency-Inverse Document Frequency) [20] โดย term Frequency (TF) คือการวัดความถี่ของคำที่ปรากฏในเอกสาร และ inverse document frequency (IDF) คือการวัดความสำคัญของคำในเอกสาร โดยคำที่ปรากฏบ่อยครั้งในเอกสารจะมีค่า IDF ต่ำ ตรงกันข้ามกับคำที่ปรากฏในเอกสารน้อยครั้งจะมีค่า IDF สูง ดังนั้นการคูณค่า TF และ IDF จะสามารถกำหนดได้ว่าแต่ละคำในข้อความมีความสำคัญมากน้อยเพียงใดต่อข้อความนั้น ๆ ซึ่งจะช่วยในการวิเคราะห์ข้อความที่มีความยาวและความแตกต่างกันเป็นไปได้อย่างมีประสิทธิภาพ

กระบวนการแปลงข้อความเป็นเวกเตอร์เชิงตัวเลขด้วยการคำนวณค่า TF-IDF ช่วยวิเคราะห์ความสำคัญของคำในข้อความโดยความถี่ของคำหรือ TF และความถี่ผกผันของเอกสารหรือ IDF คำนวณได้จากสมการที่ 2 และ 3

$$TF(t, d) = \frac{\sum_{t' \in d} t, d}{\sum_{t' \in d} t', d} \quad (2)$$

โดย $\sum_{t' \in d} t, d$ คือจำนวนครั้งที่คำ t ปรากฏในเอกสาร d และตัวส่วนเป็นจำนวนรวมของคำทั้งหมดในเอกสาร d

$$IDF(t) = \log \left(\frac{N}{1 + DF(t)} \right) \quad (3)$$

โดย N คือจำนวนเอกสารทั้งหมด และ $DF(t)$ คือจำนวนเอกสารที่มีคำ t ปรากฏอยู่ ซึ่งการบวก 1 ในตัวส่วนช่วยป้องกันปัญหาการหารด้วย 0

ดังนั้น TF-IDF เป็นค่าที่ช่วยถ่วงน้ำหนักของคำ โดยพิจารณาทั้งความสำคัญของคำภายในเอกสารเดียวกันและความสำคัญของคำเมื่อเปรียบเทียบกับเอกสารทั้งหมดในชุดข้อมูล ดังสมการที่ 4 ซึ่งจะช่วยระบุคำสำคัญในข้อความได้อย่างมีประสิทธิภาพ

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (4)$$

3.4.3) การวัดความคล้ายคลึงโคไซน์ (Cosine Similarity) [20] สำหรับวัดความคล้ายคลึงของเวกเตอร์ มักใช้ในงานที่เกี่ยวข้องกับข้อความและข้อมูลที่มีมิติสูง เนื่องจากไม่ถูกกระทบจากขนาดของเวกเตอร์ และคำนวณได้รวดเร็ว งานวิจัยนี้วัดความคล้ายคลึงระหว่างบทวิจารณ์แต่ละรายการกับคำสำคัญที่อยู่ในรูปแบบเวกเตอร์เชิงตัวเลข ซึ่งใช้ในการเปรียบเทียบว่าเนื้อหาของข้อความมีความคล้ายคลึงกันมากน้อยเพียงใดระหว่างบทวิจารณ์และคำสำคัญจากบทวิจารณ์การท่องเที่ยว จากสมการที่ 5

$$\text{Cosine similarity} = \frac{A \cdot B}{\|A\| \|B\|} \quad (5)$$

โดย $A \cdot B$ คือผลคูณเชิงสเกลาร์ (dot product) ระหว่างเวกเตอร์ A และ B และ $\|A\| \|B\|$ คือขนาดของเวกเตอร์ A และ B ตามลำดับ

จากสมการที่ 5 ค่าความคล้ายคลึงโคไซน์ จะมีค่าระหว่าง -1 ถึง 1 โดยค่าเข้าใกล้กับ 1 ยิ่งมากแสดงว่าเวกเตอร์ทั้งสองมีทิศทางเดียวกัน หรือมีความคล้ายคลึงกันสูง ในทางกลับกัน หากค่าเข้าใกล้กับ 0 แสดงว่าเวกเตอร์ทั้งสองไม่มีความคล้ายคลึงกัน และหากค่าเข้าใกล้กับ -1 แสดงว่าเวกเตอร์ทั้งสองมีทิศทางตรงกันข้าม

3.4.4) การวิเคราะห์คำสำคัญ (Keyword Analysis) จากบทวิจารณ์ของนักท่องเที่ยวต่างชาติต่ออุทยานแห่งชาติเขาใหญ่ ด้วยการประมวลผลภาษาธรรมชาติ จากการสร้างคำสำคัญจากการวิเคราะห์คำอรรถวิสัย [21] เพื่อสร้างคำสำคัญ โดยเน้นไปที่องค์ประกอบของแหล่งท่องเที่ยว 4 องค์ประกอบ ประกอบด้วย สิ่งดึงดูดใจ การเข้าถึงสถานที่ท่องเที่ยว สิ่งอำนวยความสะดวกและกิจกรรมการท่องเที่ยว เช่น คำว่า “staff and service” สำหรับองค์ประกอบด้านสิ่งอำนวยความสะดวก หรือคำว่า “hiking trail” สำหรับองค์ประกอบเกี่ยวกับกิจกรรมในอุทยาน เพื่อนำมาเปรียบเทียบความคล้ายคลึงกันระหว่างบทวิจารณ์ และคำสำคัญที่ถูกแปลงเป็นเวกเตอร์เชิงตัวเลขด้วย ค่าความคล้ายคลึงโคไซน์ โดยบทวิจารณ์ที่มีค่าความคล้ายคลึงสูงกว่าค่าที่กำหนดไว้จะถูกเลือกมาวิเคราะห์คะแนนความรู้สึก เพื่อประเมินความคิดเห็นของนักท่องเที่ยว เช่น ข้อความระบุว่า “the staff provided detailed and polite information” จะได้รับคะแนนเชิงบวก ในขณะที่ข้อความ เช่น “the hiking trail lacks sufficient signage” จะได้รับคะแนนเชิงลบ ก่อนจะนำมารายงานผลตามองค์ประกอบของแหล่งท่องเที่ยวร่วมกับการนำเสนอภาพข้อมูล

เพื่อให้เห็นภาพรวมของข้อมูลในเชิงปริมาณและเชิงคุณภาพที่สะท้อนถึงความต้องการของนักท่องเที่ยว

3.4.5) การสรุปข้อความ (Text Summarization) เป็นเทคนิคในการสรุปเนื้อหาของข้อความให้สั้นลง โดยยังคงรักษาสาระสำคัญของเนื้อหาไว้ให้ได้มากที่สุด [21] เทคนิคนี้เป็นส่วนหนึ่งของการประมวลผลภาษาธรรมชาติที่มีบทบาทสำคัญในการช่วยให้สามารถสื่อสารข้อมูลที่สำคัญไปยังผู้รับได้อย่างรวดเร็ว

ปัจจุบันการสรุปข้อความแบ่งเป็น 2 ประเภทหลัก ได้แก่ extractive summarization โดยการเลือกประโยคหรือวลีที่สำคัญจากต้นฉบับมาใช้โดยตรง จึงมีข้อดีในด้านความแม่นยำเนื่องจากสามารถสรุปโดยยังคงความถูกต้องและตรงกับเนื้อหาดั้งเดิม แต่ยังมีข้อจำกัดในเรื่องของการกระชับข้อมูล รวมถึงอาจทำให้ข้อความบางส่วนที่มีความสำคัญถูกกลืนหายไป และ abstractive summarization วิธีนี้ใช้โมเดลปัญญาประดิษฐ์ในการทำความเข้าใจบริบทของข้อความต้นฉบับและเขียนสรุปขึ้นมาใหม่ในรูปแบบที่กระชับและเป็นธรรมชาติมากขึ้น ซึ่งช่วยให้การสรุปมีความเป็นธรรมชาติ และง่ายต่อการทำความเข้าใจ เนื่องจากโมเดลสามารถเข้าใจบริบทของข้อความต้นฉบับและเลือกคำที่เหมาะสมเพื่อสร้างข้อความที่มีความหมายใหม่ได้อย่างไรก็ตาม การสร้างข้อความใหม่อาจทำให้เกิดการเข้าใจผิดหรือการสูญเสียข้อมูลสำคัญไป นอกจากนี้ยังต้องการทรัพยากรในการคำนวณที่สูงและมีความซับซ้อนในการฝึกโมเดล

จะเห็นได้ว่าทั้งสองวิธีต่างมีจุดเด่นและข้อจำกัดที่แตกต่างกัน การเลือกใช้จึงควรพิจารณาตามลักษณะของข้อมูลและวัตถุประสงค์การใช้งาน งานวิจัยนี้พิจารณาเลือกประโยคหรือวลีที่สำคัญจากต้นฉบับมาใช้โดยตรงจากการจับคู่ความสัมพันธ์ระหว่างเวกเตอร์ก่อนจะปรับแต่งกระบวนการสรุปความสำหรับใช้ในเนื้อหาที่เฉพาะเจาะจงสำหรับการท่องเที่ยวและนำเสนอในรูปแบบความเรียงประกอบภาพข้อมูลต่อไป

4) ผลการวิจัยและอภิปรายผล

งานวิจัยนี้พัฒนาขึ้นการเขียนโปรแกรมภาษาไพธอน (python) บน Google Colaboratory แบ่งผลการวิจัยและอภิปรายผลออกเป็น 2 ส่วนตามวัตถุประสงค์ ได้แก่ เพื่อวิเคราะห์บทวิจารณ์ด้วยเครื่องมือวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรม และเพื่อสร้างรายงานเชิงวิเคราะห์บทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติต่ออุทยานแห่งชาติเขาใหญ่ มีรายละเอียดดังนี้

4.1) ผลการวิเคราะห์บทวิจารณ์ด้วยเครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรม

การวิเคราะห์ความรู้สึกจากข้อความโดยอาศัยพจนานุกรมที่มีการสร้างคำศัพท์และกำหนดคะแนนความรู้สึกไว้ล่วงหน้าแบ่งผลการวิจัยและอภิปรายผลเป็น 5 ขั้นตอน ได้แก่

4.1.1) ผลการรวบรวมข้อมูลบทวิจารณ์สาธารณะ ด้วยการสกัดและดึงข้อมูลจากเว็บไซต์ที่เปิดให้นักท่องเที่ยวเขียนบทวิจารณ์สาธารณะต่ออุทยานแห่งชาติเขาใหญ่ตั้งแต่วันที่ 1 มกราคม พ.ศ. 2563 ถึง 31 ธันวาคม พ.ศ. 2567 รวมระยะเวลาทั้งสิ้น 4 ปี ซึ่งถือว่าสะท้อนภาพรวมของแหล่งท่องเที่ยวในปัจจุบัน และเหมาะสมต่อการนำมาวิเคราะห์เพื่อเป็นแนวทางพัฒนาศักยภาพของแหล่งท่องเที่ยว

งานวิจัยนี้รวบรวมข้อมูลด้วยการสกัดและดึงข้อมูลจากเว็บไซต์อัตโนมัติ เริ่มจากการส่งคำขอไปยังเว็บไซต์เป้าหมาย โดยใช้โปรโตคอล HTTP เพื่อดึงข้อมูลหน้าเว็บออกมาในรูปแบบของโครงสร้าง HTML หลังจากนั้น จะทำการวิเคราะห์และสกัดข้อมูลที่ต้องการผ่านโครงสร้างของ HTML เช่น แท็ก <div>, หรือ <table> ซึ่งเป็นส่วนที่เก็บข้อมูลหลักของเว็บไซต์ เพื่อเก็บรวบรวมข้อมูล ประกอบด้วยวันที่ ค่าคะแนนความพึงพอใจ และบทวิจารณ์สาธารณะที่สามารถนำมาใช้ในการวิเคราะห์จุดเด่นและข้อเสนอแนะของอุทยานแห่งชาติเขาใหญ่ได้

4.1.2) ผลการประมวลผลข้อมูลเบื้องต้น งานวิจัยนี้มุ่งเน้นการวิเคราะห์ข้อความจากนักท่องเที่ยวชาวต่างชาติ โดยคัดกรองเฉพาะบทวิจารณ์ที่เขียนด้วยภาษาอังกฤษ ประกอบด้วย บทวิจารณ์และคะแนนความพึงพอใจ จำนวน 12,035 รายการ ผ่านการทำความสะอาดข้อมูล โดยมีวัตถุประสงค์เพื่อปรับปรุงคุณภาพของข้อมูลให้มีความถูกต้อง สมบูรณ์ และพร้อมสำหรับการนำไปใช้งาน เริ่มต้นจากการระบุและจัดการคำที่หายไป ซึ่งอาจเกิดจากการบันทึกข้อมูลที่ไม่สมบูรณ์หรือข้อผิดพลาดในการรวบรวมข้อมูล โดยหากพบรายการใดที่มีค่าคะแนนความพึงพอใจหายไป จะพิจารณาลบแถวที่มีค่าหายไป เนื่องจากไม่ส่งผลกระทบต่อความสมบูรณ์ของข้อมูล นอกจากนี้ ยังมีการจัดรูปแบบข้อมูลให้เป็นมาตรฐาน เช่น การแปลงค่าข้อความให้เป็นตัวอักษรพิมพ์เล็กทั้งหมด และการลบอักขระพิเศษที่ไม่จำเป็น เพื่อให้ข้อมูลมีคุณภาพและความน่าเชื่อถือก่อนนำไปใช้

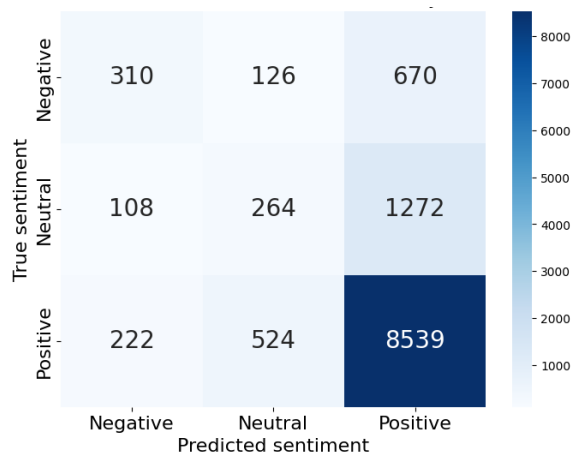
4.1.3) ผลการประเมินผลเครื่องมือสำหรับวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรมด้วยค่าความถูกต้องที่คำนวณจากสัดส่วนของการจำแนกข้อมูลว่าเป็นเชิงบวก เชิงลบ

หรือเป็นกลาง เมื่อเทียบกับผลจริงที่ได้จากการให้คะแนนความพึงพอใจ ได้แก่ TextBlob, Flair และ VADER จากการวัดซ้ำจนครบ 3 รอบ แล้วนำผลการประเมินมาหาค่าเฉลี่ย ดังตารางที่ 3

ตารางที่ 3 : ค่าความถูกต้องเฉลี่ยจากการวิเคราะห์ความรู้สึกจากข้อความ

sentiment analysis tool	accuracy
TextBlob	0.7198
Flair	0.7324
VADER	0.7575

จากตารางที่ 3 ค่าความถูกต้องเฉลี่ยจากการวิเคราะห์ความรู้สึกจากข้อความระหว่าง TextBlob, Flair และ VADER มีค่าความถูกต้องเฉลี่ยใกล้เคียงกันคิดเป็นร้อยละ 72, 73 และ 76 ตามลำดับโดย VADER มีค่าความถูกต้องเฉลี่ยสูงที่สุด



รูปที่ 2 : เมทริกซ์ความสัมพันธ์การจำแนกความรู้สึกจากข้อความด้วย VADER

เมื่อพิจารณาจากเมทริกซ์ความสัมพันธ์ (confusion matrix) ซึ่งใช้ประเมินประสิทธิภาพของการจำแนกความรู้สึกเป็นเชิงบวก เชิงลบ และเป็นกลางของ VADER ดังรูปที่ 2 พบว่า การจำแนกความรู้สึกจากข้อความได้อย่างถูกต้องประมาณ 3 ใน 4 ของทั้งหมด เมื่อพิจารณาในรายละเอียด พบว่า สามารถจำแนกข้อความเชิงบวกได้ดีมาก โดยมีจำนวนตัวอย่างที่ทำนายถูกต้องถึง 8,539 ตัวอย่าง แต่ยังมีข้อผิดพลาดในการจำแนกข้อความที่เป็นกลางและข้อความเชิงลบอยู่ในระดับที่ค่อนข้างสูง โดยจำแนกความรู้สึกผิดว่าเป็นเชิงบวก ทั้งที่เป็นเชิงลบหรือเป็นกลางรวมกันถึง 746 ตัวอย่าง แสดงให้เห็นถึงแนวโน้มที่ให้น้ำหนักกับข้อความเชิงบวกมากเกินไป (bias toward positive) เช่นเดียวกับการจำแนกข้อความที่เป็นกลาง ซึ่งทำนายผิดไปเป็นเชิงบวกจำนวน

1,272 ตัวอย่าง ถือเป็นอัตราส่วนที่สูงมากเมื่อเทียบกับข้อความที่เป็นกลางทั้งหมด นอกจากนี้ ยังมีข้อผิดพลาดในการจำแนกข้อความเชิงลบบางส่วนที่ถูกจำแนกผิดเป็นเชิงบวกหรือเป็นกลางรวมกันถึง 796 ตัวอย่าง

อย่างไรก็ดี นอกเหนือจากค่าความถูกต้องที่เป็นตัวชี้วัดประสิทธิภาพในการทำงานแล้ว เครื่องมือสำหรับวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรมแต่ละแบบ ยังมีจุดเด่นและข้อจำกัดในการทำงานแตกต่างกันไป โดย TextBlob เหมาะสำหรับการใช้งานทั่วไปที่ต้องการความง่ายในการใช้งานส่วน Flair เหมาะสำหรับงานที่ต้องการเข้าใจบริบทของข้อความ แต่อาจต้องแลกกับการใช้ทรัพยากรการประมวลผลมากขึ้น ในทางกลับกัน VADER นั้นเหมาะสำหรับงานที่ต้องการความเร็วในการวิเคราะห์ข้อมูลและรองรับข้อความภาษาธรรมชาติจากสื่อสังคมออนไลน์ ดังนั้นการเลือกใช้งานควรพิจารณาจากลักษณะของข้อมูล และข้อจำกัดด้านทรัพยากรในการประมวลผลเป็นปัจจัยร่วมในการพิจารณาอีกด้วย

งานวิจัยนี้พิจารณาเลือก VADER ที่เป็นเครื่องมือสำหรับวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรมที่ออกแบบมาให้ใช้งานกับข้อความที่ไม่มีโครงสร้างได้ทันที โดยไม่ต้องผ่านกระบวนการตัดคำ การกำจัดคำหยุด และสามารถวิเคราะห์ข้อมูลจากข้อความทั้งหมดได้ทันที โดยไม่จำเป็นต้องแปลงเป็นเวกเตอร์ตัวเลขก่อน นอกจากนี้ยังมีค่าความถูกต้องเฉลี่ยสูงที่สุดแล้วยังเหมาะสำหรับงานที่ต้องการความเร็วในการวิเคราะห์ข้อความภาษาธรรมชาติบนอุปกรณ์ที่มีข้อจำกัดด้านทรัพยากรอย่างสมาร์ทโฟนได้อย่างมีประสิทธิภาพ ซึ่งเป็นปัจจัยสำคัญที่ควรพิจารณาสำหรับการใช้งานในสถานการณ์จริง

4.1.4) ผลการวิเคราะห์ความเอนเอียงในบทวิจารณ์ เพื่อประเมินว่าการวิเคราะห์ความรู้สึกจากบทวิจารณ์นั้นสามารถสะท้อนความคิดเห็นของนักท่องเที่ยวได้อย่างถูกต้องหรือไม่ โดยเฉพาะในกรณีที่คะแนนความพึงพอใจ และเนื้อหาของบทวิจารณ์มีความขัดแย้งกัน โดยนำบทวิจารณ์จำนวน 12,035 รายการที่ผ่านการวิเคราะห์ความรู้สึกจาก VADER มาพิจารณาบทวิจารณ์ที่มีความเอนเอียง เพื่อทำความเข้าใจทิศทางและแนวโน้มของบทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติในแง่บวก แง่ลบ หรือเป็นกลาง โดยพิจารณาบทวิจารณ์ที่มีความเอนเอียงในทางบวกหรือลบอย่างชัดเจน พบว่า บทวิจารณ์ที่คะแนนเท่ากับ 5 แต่มีความรู้สึกเป็นลบ มีจำนวนเท่ากับ 2,912 รายการ ในทางกลับกันบทวิจารณ์ที่คะแนนเท่ากับ 1 แต่มีความรู้สึกเป็นบวก มีจำนวน

เท่ากับ 448 รายการ เมื่อเรียงลำดับตามค่าคะแนนความรู้สึกจากมากไปน้อย รายละเอียดดังตารางที่ 4 และ 5

ตารางที่ 4 : ตัวอย่างบทวิจารณ์ที่มีค่าคะแนนความพึงพอใจเท่ากับ 5
แต่ค่าความรู้สึกเป็นลบ (negative)

review	sentiment score
it is worth the effort to get there. a treat to see wildlife in its natural unspoiled environment. used to see wildlife at home, but this topped it to see snakes, scorpions, millions of bats, deer, giant squirrels, gibbons, macaques, and even the only croc in the park near the waterfall trail, all in one day. only negative is the late morning opening and late afternoon closing of the park. probably why we missed the elephants.	-0.7263
a must visit if you're in khao yai being city dweller, 1st time exploring and experiencing the nature and didn't even bother about being sucked dry by leeches. encountered a couple of wildlife and the view on the top is to die for.	-0.6194
...	...
our best experience in Thailand. mostly visited by Thai tourists, and not really by foreigners it seems. it's a real wild jungle, with many animals to see crocodiles, monkeys, spiders, snakes. for a real jungle experience, go visit it with a ranger and stay there overnight in a tent. but closely follow the rangers or guides advice if you don't want to be attacked by leeches.	-0.1343

จากตารางที่ 4 บทวิจารณ์จากนักท่องเที่ยวต่างชาติที่มีค่าคะแนนความพึงพอใจเท่ากับ 5 แต่ค่าความรู้สึกเป็นลบในหลายแง่มุม ซึ่งสามารถจำแนกได้เป็นประเด็นหลัก ๆ ได้แก่ ค่าธรรมเนียมที่ไม่เท่าเทียม ข้อจำกัดในการเข้าถึง ปัญหาด้านการบริการ และความคาดหวังที่ไม่ตรงกับความเป็นจริง โดยเฉพาะการเปรียบเทียบค่าใช้จ่ายระหว่างนักท่องเที่ยวชาวไทยและชาวต่างชาติที่ต่างกันอย่างชัดเจน ประเด็นต่อมาคือข้อจำกัดในการเข้าถึงและกิจกรรมต่าง ๆ ที่ไม่ตอบสนองต่อความ

ต้องการของผู้เยี่ยมชมอย่างเต็มที่ อีกทั้งการจำกัดเวลาเปิดและปิดอุทยานทำให้พลาดโอกาสในการชมสัตว์ป่าในช่วงเวลาที่แตกต่างกันตามแต่ละฤดูกาล นอกจากนี้ ยังมีการกล่าวถึงการบริการที่ไม่เป็นไปตามความคาดหวัง บางบทวิจารณ์ได้ระบุว่าผู้นำเที่ยวขาดความพร้อมในการให้บริการ เช่น ไม่แจ้งข้อมูลหรือคำแนะนำที่ชัดเจนในเรื่องความปลอดภัย หรือคำแนะนำต่อการปรับตัวตามสภาพอากาศและสถานการณ์ต่าง ๆ ที่เหมาะสมจึงทำให้รู้สึกที่ไม่คุ้มค่ากับราคา อีกทั้งพลาดการเห็นสัตว์ที่มีชื่อเสียงอย่างช้าง หรือการที่สภาพเส้นทางเดินไม่เป็นไปตามคำแนะนำ อาจนำไปสู่ความผิดหวังนักท่องเที่ยว แม้สถานที่นั้นจะมีความงดงาม จึงแสดงความรู้สึกเชิงลบแม้ว่าจะได้รับประสบการณ์ที่ดีในบางแง่มุมจากการท่องเที่ยว

ตารางที่ 5 : ตัวอย่างบทวิจารณ์ที่มีค่าคะแนนความพึงพอใจเท่ากับ 1
แต่ค่าความรู้สึกเป็นบวก (positive)

review	sentiment score
we were excited to visit but were put off right away by the \$400 baht charge per adult for foreigners. for Thai people, it's only \$40 baht, so visitors are charged 10 times the price. they also charge extra based on the vehicle you're driving. crazy. besides the price, the main road going through the park is more like a speedway for people crossing from one side to the other.	0.9536
it's kind of like a cheap toll road for Thais. if you've ever visited a national park in other countries, it's a different experience altogether. we weren't expecting to see animals but were hopeful. we saw a few monkeys on the road and one deer. that's it. it's a scenic park to drive through, but literally, that's all we did, drive through. i wouldn't recommend coming here. save your \$400 baht and have a nice meal out instead.	0.9536

ตารางที่ 5 : ตัวอย่างบทวิจารณ์ที่มีค่าคะแนนความพึงพอใจเท่ากับ 1 แต่ค่าความรู้สึกเป็นบวก (positive) (ต่อ)

Review	sentiment score
did not bother, fee too high for foreigners, 400 baht. for locals it's 30 baht, why the drastic difference. what a turn off for foreigners. we have visited great national parks all over the world which have reasonable fees or are free.	0.8568
...	...
i do not recommend the night tour offered from the park, since the guide spoke no English (not even the name of the animals), and showed no interest in looking for animals. just saw deer.	0.1002

เช่นเดียวกับ บทวิจารณ์ที่มีค่าคะแนนความพึงพอใจเท่ากับ 1 แต่ค่าความรู้สึกเป็นบวกจากตารางที่ 5 พบว่า หลายบทวิจารณ์แสดงความผิดหวังจากปัญหาด้านค่าธรรมเนียมที่ไม่เท่าเทียม การบริการที่ไม่ได้มาตรฐาน และประสบการณ์จริงที่ไม่ตรงกับ ความคาดหวัง แม้จะพึงพอใจในแง่ของความงดงามของธรรมชาติ และสัตว์ป่า แต่ปัญหาด้านการจัดการ และค่าใช้จ่ายที่สูงเกินไป กลับสร้างความผิดหวังและส่งผลกระทบต่อคะแนนความพึงพอใจโดยรวมที่ลดลง

4.2) ผลการสร้างรายงานเชิงวิเคราะห์

รายงานเชิงวิเคราะห์จากบทวิจารณ์ของนักท่องเที่ยวชาวต่างชาติต่ออุทยานแห่งชาติเขาใหญ่ที่ระบุจุดเด่น รวมถึงข้อจำกัดของการท่องเที่ยว อันจะเป็นแนวทางในการพัฒนาและยกระดับแหล่งท่องเที่ยวให้สอดคล้องกับมาตรฐานสากล แบ่งผลการวิจัยและอภิปรายผลออกเป็น 5 ขั้นตอน ประกอบด้วย

4.2.1) ผลการลบคำหยาบ เพื่อกำจัดข้อมูลที่ไม่จำเป็น และปรับข้อความให้อยู่ในรูปแบบที่เหมาะสมกับการประมวลผล ประกอบด้วยการแปลงข้อความเป็นตัวอักษรพิมพ์เล็กทั้งหมด การแปลงคำให้อยู่ในรูปแบบพื้นฐาน รวมถึงการลบอักขระพิเศษ และเครื่องหมายวรรคตอนที่ไม่มีความหมายของข้อความ ดังตารางที่ 6 คำหยาบที่ถูกลบออก ประกอบด้วย “a”, “an”, “and”, “are”, “as”, “at”, “be”, “but”, “by”, “can”, “do”,

“for”, “from”, “has”, “have”, “if”, “in”, “is”, “it”, “its”, “not”, “of”, “on”, “or”, “so”, “that”, “the”, “this”, “to”, “was”, “were”, “will”, “with”, “you” และ “your”

ตารางที่ 6 : ตัวอย่างบทวิจารณ์ที่ผ่านกระบวนการประมวลผลข้อมูลเบื้องต้น

No.	Review	Rating
1.	nice park many trail waterfall hike trail without guide map available visitor center people complain pay foreigner see difference salary foreigner earn much common thai people support higher entrance fee conserve park	5
2.	incredibly big scenic park local pay 40 thb per head foreigner pay 100 thb well worth visit wild animal roam massive park go winter month weather cool windy need car go entire park	5
3.	expect much see animal may find elephant solid waste go jungle trekking without get tired national park route not long almost no climbing take 2 hr 3.3 km hundred year old tree make noise place frighten animal throw away garbage feed animal bring water snack	4
...	...	...
12035.	park lovely want go camping forget despite numerous campsite mark map one genuine one like supermarket car park	1

จากตารางที่ 6 ข้อความที่ผ่านการลบคำหยาบ และผ่านการลบอักขระ เช่น ., “, ”, !, ? และ , จากนั้นแปลงคำให้อยู่ในรูปแบบพื้นฐาน ได้แก่ “trails” เป็น “trail”, “waterfalls” เป็น “waterfall”, “hiking” เป็น “hike”, “guides” เป็น “guide”, “maps” เป็น “map”, “visitors” เป็น “visitor”, “complaining” เป็น “complain”, “paying” เป็น “pay”, “foreigners” เป็น “foreigner”, “earning” เป็น “earn”, “supports” เป็น “support”, “conserving” เป็น “conserve”, “pays” เป็น “pay”, “heads” เป็น “head”, “roaming” เป็น “roam”, “going” เป็น “go”, “expecting” เป็น “expect”, “seeing” เป็น “see”, “animals” เป็น “animal”, “finding” เป็น “find”, “takes” เป็น “take”, “years” เป็น “year”, “trees” เป็น “tree”, “making” เป็น “make”, “frightening” เป็น

“frighten”, “throwing” เป็น “throw”, “feeding” เป็น “feed”, “bringing” เป็น “bring”, “wanting” เป็น “want”, “camping” เป็น “camp”, “forgetting” เป็น “forget”, “marking” เป็น “mark” และ “likes” เป็น “like”

4.2.2) ผลการแปลงข้อความเป็นเวกเตอร์เชิงตัวเลขด้วย TF-IDF ซึ่งช่วยกำหนดค่าความสำคัญของแต่ละคำในเอกสารได้ โดยคำที่ปรากฏบ่อยในข้อความเดียวกันจะได้รับค่าน้ำหนักที่ต่ำลง ในขณะที่คำที่มีความสำคัญต่อเนื้อหาจะได้รับค่าน้ำหนักที่สูงขึ้น โดยเรียกใช้ TfidfVectorizer เพื่อคำนวณค่าความสำคัญของคำในประโยคตัวอย่าง 3 ประโยค ประกอบด้วย “the park is good”, “the park is bad” และ “awesome is the awesome” แทนด้วย doc1, doc2 และ doc3 ตามลำดับ รายละเอียดดังตารางที่ 7 ถึง 9

ตารางที่ 7 : การหาความถี่ของคำแต่ละคำในเอกสาร

TF (term-frequency)						
	the	park	is	good	bad	awesome
doc1	1/4 =0.25	1/4 =0.25	1/4 =0.25	1/4 =0.25	0	0
doc2	1/4 =0.25	1/4 =0.25	1/4 =0.25	0	1/4 =0.25	0
doc3	1/4 =0.25	0	1/4 =0.25	0	0	2/4 =0.50

ตารางที่ 8 : การคำนวณค่าความสำคัญของคำ

IDF (inverse document-frequency)						
	the	park	is	good	bad	awesome
value	$\log(3/3)$ =0	$\log(3/2)$ =0.1761	$\log(3/3)$ =0	$\log(3/1)$ =0.4771	$\log(3/1)$ =0.4771	$\log(3/1)$ =0.4771

ตารางที่ 9 : ผลการแปลงข้อความเป็นเวกเตอร์คุณลักษณะ

feature vector						
	the	park	is	good	bad	awesome
doc1	0	0.0440	0	0.1193	0	0
doc2	0	0.0440	0	0	0.1193	0
doc3	0	0	0	0	0	0.2386

จากตารางที่ 7 การคำนวณค่าความถี่ของคำจากจำนวนครั้งที่คำปรากฏในเอกสาร ทารด้วยจำนวนคำทั้งหมดในเอกสารนั้น คำที่ได้แสดงความถี่ของคำแต่ละคำในเอกสารหนึ่ง ๆ เพื่อคำนวณค่าความสำคัญของคำ ดังตารางที่ 8 คำที่มีค่าสูงกว่า

แสดงว่าคำปรากฏในเอกสารน้อยกว่า และมีความสำคัญมากกว่า และตารางที่ 9 เวกเตอร์คุณลักษณะ (feature vector) ที่ได้จากการคำนวณ TF-IDF จะแสดงถึงความสำคัญของคำในเอกสารแต่ละฉบับที่มีความแตกต่างกัน เช่น “good”, “bad” และ “awesome” ที่ปรากฏในเอกสาร doc1, doc2 และ doc3 จึงสามารถใช้ในการแยกแยะคำสำคัญได้ดีขึ้น ในทางกลับกัน “the” และ “is” มีค่า TF-IDF ต่ำหรืออาจเท่ากับ 0 แสดงถึงเป็นคำที่ปรากฏบ่อย แต่ไม่ค่อยมีความสำคัญในเอกสาร

4.2.3) ผลการวัดความคล้ายคลึงโคไซน์ เป็นวิธีการวัดความคล้ายคลึงระหว่างเวกเตอร์ของสองประโยคหรือเอกสาร โดยคำนวณจากผลคูณจุดของเวกเตอร์ทั้งสองและย่อด้วยความยาวของเวกเตอร์แต่ละอัน

เมื่อข้อความถูกแปลงเป็นเวกเตอร์เชิงตัวเลขด้วย TF-IDF จะนำมาคำนวณผลคูณจุด (dot product) ระหว่างเวกเตอร์แต่ละคู่ประโยค โดยนำค่าของ TF-IDF จากเอกสารแต่ละคำมาคูณกันแล้วหาผลรวม หากเอกสารมีค่าเหมือนกันมาก ค่าจะสูงขึ้น ดังตารางที่ 10 จากนั้นจะคำนวณความยาว (norm) ของเวกเตอร์ เพื่อแสดงขนาดของเวกเตอร์ หากค่า TF-IDF ยิ่งสูงแสดงว่าเอกสารมีค่าเฉพาะมากขึ้น ดังตารางที่ 11 ก่อนจะคำนวณค่าความคล้ายคลึงโคไซน์สำหรับแต่ละคู่ของประโยค

ตารางที่ 10 : การคำนวณผลคูณจุดระหว่างเวกเตอร์แต่ละคู่ประโยค

pair	calculation	dot product
doc1 * doc2	$(0 \times 0) + (0.044 \times 0.044) + (0 \times 0) + (0.119 \times 0) + (0 \times 0.119) + (0 \times 0)$	0.001936
doc1 * doc3	$(0 \times 0) + (0.044 \times 0) + (0 \times 0) + (0.119 \times 0) + (0 \times 0) + (0 \times 0.238)$	0
doc2 * doc3	$(0 \times 0) + (0.044 \times 0) + (0 \times 0) + (0 \times 0) + (0.119 \times 0) + (0 \times 0.238)$	0

ตารางที่ 11 : การคำนวณความยาวของเวกเตอร์

document	calculation	norm
doc1	$\sqrt{(0^2 + 0.044^2 + 0^2 + 0.119^2 + 0^2 + 0^2)}$	0.1268
doc2	$\sqrt{(0^2 + 0.044^2 + 0^2 + 0^2 + 0.119^2 + 0^2)}$	0.1268
doc3	$\sqrt{(0^2 + 0^2 + 0^2 + 0^2 + 0^2 + 0.238^2)}$	0.238

ตารางที่ 12 : การคำนวณค่าความคล้ายคลึงโคไซน์

pair	calculation	cosine similarity
doc1 , doc2	$\frac{0.001936}{0.1268 \times 0.1268} = \frac{0.001936}{0.016097}$	0.1204
doc1 , doc3	$\frac{0}{0.1268 \times 0.2380} = \frac{0}{0.0302}$	0
doc2 , doc3	$\frac{0}{0.1268 \times 0.2380} = \frac{0}{0.0302}$	0

ดังตารางที่ 12 ค่าความคล้ายคลึงระหว่าง doc1 กับ doc2 เท่ากับ 0.1204 แสดงถึงความคล้ายคลึงกันเล็กน้อย เนื่องจากมีคำที่ใช้เหมือนกันหลายคำ ในทางตรงกันข้าม ความคล้ายคลึงระหว่าง doc1 กับ doc3 และ doc2 กับ doc3 เท่ากับ 0 เนื่องจากไม่มีคำใดเลยที่ใช้ร่วมกัน

4.2.4) ผลการวิเคราะห์คำสำคัญ งานวิจัยนี้สร้างคำสำคัญพิจารณาจากการวิเคราะห์คำอัตวิสัย (subjectivity analysis) ของบทวิจารณ์ด้วย TextBlob ที่เป็นหนึ่งในเครื่องมือสำหรับวิเคราะห์ความรู้สึกจากข้อความแบบใช้พจนานุกรมที่มีการระบุคำอัตวิสัยไว้ล่วงหน้า บนแนวคิดของข้อความที่เป็นอัตวิสัย (subjective) มักจะเป็นข้อความที่สะท้อนถึงความคิดเห็น ความรู้สึก หรือการตีความส่วนบุคคล ในทางตรงกันข้าม ข้อความที่เป็นภาววิสัย (objective) จะเป็นข้อความที่เกี่ยวข้องกับข้อมูลที่เป็นข้อเท็จจริง หรือเป็นคำอธิบายที่ไม่เกี่ยวข้องกับความรู้สึกหรือความคิดเห็นส่วนตัว ดังสมการที่ 6

$$Subjectivity = \frac{\sum (subjectivity \text{ of each word})}{number \text{ of word with subjectivity value}} \quad (6)$$

จากสมการที่ 6 คำอัตวิสัยจะอยู่ในช่วงตั้งแต่ 0 ถึง 1 โดยค่า 0 หมายถึงเป็นข้อความที่มีความเป็นกลางสูง ในทางกลับกัน ค่า 1 หมายถึงข้อความที่เต็มไปด้วยความคิดเห็นหรือความรู้สึกส่วนบุคคล และค่ากลางระหว่าง 0 และ 1 แสดงถึงระดับของความเป็นส่วนตัวที่มีความหลากหลายจากข้อความที่เป็นกลางไปสู่ข้อความที่มีความเป็นส่วนตัวสูง ตัวอย่างข้อความคือ “nice park with many trails and waterfall”

ตารางที่ 13 : ตารางตัวอย่างคำอัตวิสัยของคำแต่ละคำในคลังคำศัพท์

subjectivity scores of each word							
word	nice	park	with	many	trails	and	waterfall
subjectivity	1.0	0	0	0.5	0	0	0

จากตารางที่ 13 คำอัตวิสัยที่บ่งบอกถึงค่าที่เป็นความคิดเห็นส่วนบุคคลที่แสดงออกผ่านคำหรือกลุ่มคำในข้อความ โดยคำว่า “nice” มีคำอัตวิสัยเท่ากับ 1.0 สะท้อนให้เห็นว่าคำนี้เป็นความคิดเห็นส่วนบุคคลต่อมุมมองเชิงบวกอย่างเด่นชัด เช่นเดียวกับคำว่า “many” มีคำอัตวิสัยเท่ากับ 0.5 ซึ่งบ่งบอกถึงค่าที่เป็นกลางระหว่างข้อเท็จจริงและความคิดเห็นส่วนบุคคล ต่อลักษณะการประเมินที่เป็นเชิงปริมาณ แต่ในขณะเดียวกันก็สะท้อนถึงการยอมรับหรือความพึงพอใจในความหลากหลายหรือจำนวนที่มากในทางตรงกันข้าม คำว่า “park”, “trails” และ “waterfall” ที่มีคำอัตวิสัยเท่ากับ 0.0 ที่บ่งชี้ถึงความเป็นข้อเท็จจริง และไม่มีความคิดเห็นส่วนตัวแฝงอยู่ ดังนั้นคำเหล่านี้ทำหน้าที่บ่งบอกลักษณะเฉพาะของสถานที่หรือกิจกรรม โดยไม่เพิ่มความรู้สึกเข้าไปในข้อความ เมื่อนำมาคำนวณคำอัตวิสัยดังตารางที่ 14 จะเห็นว่าค่าที่มีผลต่อคำอัตวิสัย คือ “nice” และ “many” เท่านั้น ในขณะที่คำอื่น ๆ เป็นข้อเท็จจริง ทำให้คำอัตวิสัย โดยรวมอยู่ที่ 0.75 ดังตารางที่ 15

ตารางที่ 14 : ตัวอย่างการคำนวณคำอัตวิสัย

subjectivity calculation
$Subjectivity = \frac{(1.00 + 0.50)}{2} = \frac{1.50}{2} = 0.75$

ตารางที่ 15 : ผลลัพธ์จากการคำนวณคำอัตวิสัย และความระดับความเข้มข้น

subjectivity	nice	park	with	many	trails	and	waterfall
0.75	1.0	0	0	0.5	0	0	0

จากตารางที่ 15 งานวิจัยนี้คัดเลือกคำสำคัญจากคำที่มีคำอัตวิสัยเท่ากับ 0.0 เท่านั้น จากการพิจารณาความสอดคล้องกับองค์ประกอบของแหล่งท่องเที่ยว 4 องค์ประกอบ ได้แก่ สิ่งดึงดูดใจ การเข้าถึงสถานที่ท่องเที่ยว สิ่งอำนวยความสะดวก และกิจกรรมการท่องเที่ยว องค์ประกอบละ 5 คำ เช่น “waterfall” ซึ่งเป็นหนึ่งในคำสำคัญด้านสิ่งดึงดูดใจ ก่อนจะถูกแปลงให้อยู่ในรูปแบบเวกเตอร์คำสำคัญ (keyword vector) ด้วย TF-IDF ที่สามารถนำไปเปรียบเทียบกับเวกเตอร์บทวิจารณ์ (review vector) ด้วยคำนวณค่าความคล้ายคลึงโคไซน์ ผลลัพธ์ที่ได้จะเป็นค่าตัวเลขที่บ่งบอกระดับความใกล้เคียงระหว่างคำสำคัญกับแต่ละบทวิจารณ์ ก่อนจะจัดเรียงผลลัพธ์จากค่าที่สูงสุดไปหาค่าต่ำสุดสำหรับนำมาคัดเลือกเฉพาะบทวิจารณ์ที่มีความเกี่ยวข้องสูง โดยใช้เกณฑ์ที่กำหนด (threshold) เท่ากับ 0.5 ในการกรองข้อมูล กระบวนการนี้

ช่วยให้มั่นใจว่าข้อความที่นำมาใช้ในการสรุปจะมีเนื้อหาที่สอดคล้องกับคำค้นหาในขั้นตอนต่อไป

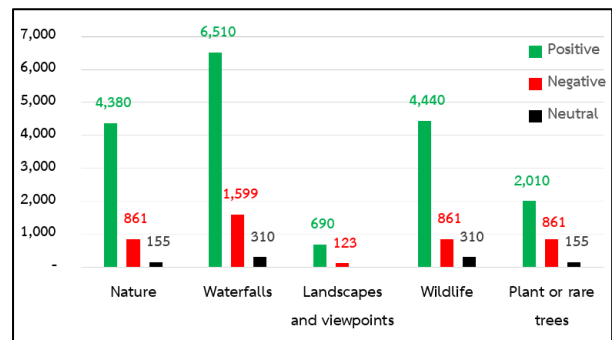
4.2.5) การสรุปข้อความที่ผ่านการคัดเลือกจะถูกปรับแต่งให้อยู่ในรูปแบบที่กระชับและเข้าใจง่ายขึ้น กระบวนการนี้ประกอบด้วยหลายขั้นตอนผสมผสานกัน ได้แก่ การกรองคำซ้ำ เช่นเดียวกับการเลือกคำสำคัญและการตัดคำที่ไม่จำเป็น ร่วมกับการลบคำที่ไม่มีความสำคัญต่อเนื้อหา จะสามารถช่วยให้ข้อความมีความกระชับขึ้นโดยไม่ทำให้สาระสำคัญเสียไป ดังตารางที่ 16 ตัวอย่างคำสำคัญคือ “waterfall”

ตารางที่ 16 : ตัวอย่างบทวิจารณ์ที่เกี่ยวข้องกับคำสำคัญ และผลการสรุปข้อความ

relative review	text summarization
My wife and i visited the national park on a small group tour to see the waterfall and to have an elephant ride, we found the park quite beautiful and a nice walk down to the waterfall, the steps are quite steep and a lot of them getting down to the viewing point for the waterfall.	Visited the National park on a small group tour to see waterfall have an elephant ride, we found quite beautiful nice walk down waterfall, steps are steep lot of them getting viewing point for waterfall.
We were at the waterfall, it's beautiful there, you can't miss it!	We were at it's there, you can't miss it! Haew Su Wat beautiful Khao Yai waterfall, adjacent Narok waterfall and it is also most famous Khao Park.
Haew su wat beautiful khao yai waterfall, adjacent to haew narok waterfall. And it is also the most famous waterfall of khao yai national park.	During rainy season, there will be water, but specialty this in summer because when water ebbs, cave or caves behind. Which can up mountain explore cave, area around covered with various trees, covering sunlight, making atmosphere inside only cool, shady comfortable.
During the rainy season, there will be a lot of water. But the specialty of this waterfall is in summer. Because when the water ebbs, it will see a small cave or caves. Behind the waterfall which can walk up the mountain to explore this small cave, the area around the waterfall is covered with various trees.	
Covering the sunlight, making the atmosphere inside only cool, shady and comfortable.	
...	
National park with a beautiful waterfall Great to see	

จากตารางที่ 16 การกรองคำซ้ำจะช่วยให้ข้อความมีความกระชับขึ้น โดยคงเนื้อหาสำคัญไว้ เช่น คำว่า “waterfall” ปรากฏหลายครั้งในบทวิจารณ์จะคงไว้เพียงครั้งเดียว เช่นเดียวกับการเลือกคำสำคัญและการตัดคำที่ไม่จำเป็น เช่น “a”, “an”, “the” หรือ “is” เป็นต้น ซึ่งเป็นคำบุพบทหรือ คำสรรพนามที่ไม่มีผลต่อความหมายหลักของข้อความ ก่อนจะจัดระเบียบข้อความใหม่ให้มีโครงสร้างที่เหมาะสม ง่ายต่อการอ่าน และสื่อความหมายได้อย่างชัดเจน โดยใช้การรวมประโยคที่มีเนื้อหาใกล้เคียงกัน หรือการปรับเปลี่ยนลำดับของข้อความเพื่อให้มีความต่อเนื่อง ก่อนจะนำไปใช้สรุปบทวิจารณ์อัตโนมัติจากข้อความจำนวนมาก ซึ่งจะช่วยให้สามารถสรุปข้อมูลที่ง่ายต่อการทำความเข้าใจของผู้ใช้ได้อย่างมีประสิทธิภาพมากขึ้น

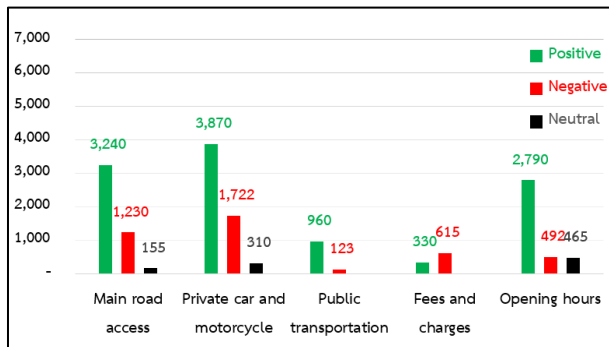
การนำเสนอข้อมูลในรูปแบบกราฟจากการนำค่าความรู้สึกจากแต่ละบทวิจารณ์ที่มีความเกี่ยวข้องมานำเสนอผลลัพธ์ในรูปแบบภาพข้อมูล จำแนกแต่ละประเด็นออกเป็นความรู้สึกเชิงบวก เชิงลบ และเป็นกลาง พร้อมกับอธิบายผลอีกครั้งใน 4 องค์ประกอบด้านการท่องเที่ยวโดยผู้วิจัย ประกอบด้วยด้านสิ่งดึงดูดใจ ด้านการเข้าถึง ด้านสิ่งอำนวยความสะดวก และด้านกิจกรรม ดังรูปที่ 4 ถึง 7



รูปที่ 4 : การวิเคราะห์คำสำคัญจากองค์ประกอบด้านสิ่งดึงดูดใจ

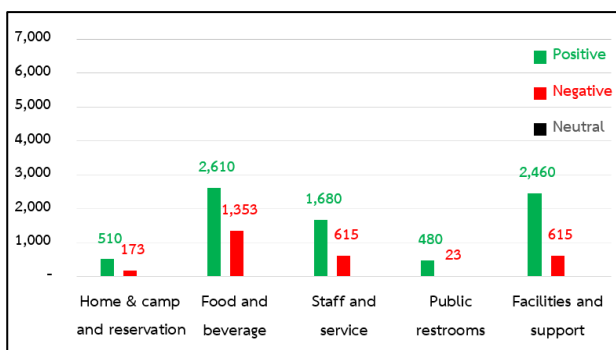
จากภาพที่ 4 อุทยานแห่งชาติเขาใหญ่ได้รับการชื่นชมอย่างสูงจากนักท่องเที่ยวในด้านสิ่งดึงดูดใจจากธรรมชาติทั้งดงและ ความหลากหลายทางชีวภาพ รวมถึงความอุดมสมบูรณ์ของพรรณไม้หายาก และสัตว์ป่านานาชนิด อีกทั้งมีน้ำตกที่สวยงามเหมาะกับการไปเยี่ยมชม อย่างไรก็ตาม บางช่วงเวลาอาจขาดแคลนน้ำส่งผลให้ไม่สามารถสัมผัสกับความงดงามของน้ำตกได้ตามที่คาดหวัง นอกจากนี้บางจุดในอุทยานที่ต้องเดินขึ้นบันไดที่ชันและจุดที่ไม่มีราวจับ ซึ่งไม่เหมาะสมกับผู้สูงอายุ หรือผู้ใช้รถเข็น ส่วนจุดชมวิวต่าง ๆ ซึ่งมีทิวทัศน์ที่หลากหลายได้รับการ

ยกย่องว่าเป็นสถานที่ที่ควรเยี่ยมชม แต่ยังมีบางส่วนที่ไม่พบเห็นสัตว์ป่าบางชนิดตามที่คาดหวัง แม้จะมีการนำเที่ยวโดยผู้เชี่ยวชาญ แต่ยังมีบางข้อเสนอแนะที่ต้องการให้การเยี่ยมชมเป็นไปตามความคาดหวังและสะดวกสบายยิ่งขึ้น



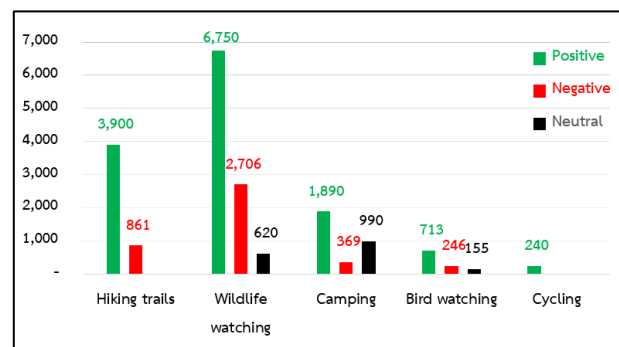
รูปที่ 5 : การวิเคราะห์ค่าสำคัญจากองค์ประกอบด้านการเข้าถึง

จากรูปที่ 5 การเข้าถึงอุทยานแห่งชาติเขาใหญ่ผ่านถนนสายหลักมีความสะดวกสบาย โดยเฉพาะเส้นทางที่เชื่อมต่อกับจุดท่องเที่ยวสำคัญต่าง ๆ เหมาะสำหรับการเดินทางด้วยรถยนต์หรือจักรยานยนต์ เนื่องจากสามารถแวะตามจุดท่องเที่ยวได้ง่าย แม้จะยังมีข้อจำกัดด้านความแออัดในช่วงฤดูการท่องเที่ยว อีกทั้งการใช้ความเร็วสูงในเขตอุทยานนั้นเสี่ยงต่อความปลอดภัยทั้งของผู้ร่วมใช้ทางและสัตว์ป่า รวมถึงการใช้รถใช้ถนนนอกเวลาทำการยังอาจรบกวนผู้พักค้างคืนและสัตว์ป่า ขณะที่บริการขนส่งสาธารณะยังมีข้อจำกัด ทำให้ต้องพึ่งพารถรับจ้างท้องถิ่นในราคาที่ไม่เป็นมาตรฐาน อย่างไรก็ตาม เวลาทำการยังเป็นข้อจำกัดในการชมสัตว์ป่าในช่วงเช้ามีดหรือช่วงค่ำที่แตกต่างกันไปในแต่ละฤดูกาล ขณะเดียวกันค่าธรรมเนียมที่แตกต่างกันระหว่างชาวไทยและชาวต่างชาติยังเป็นอีกประเด็นที่ได้รับการวิจารณ์ถึงความเหมาะสม แม้จะเป็นชาวต่างชาติที่มีบัตรอนุญาตทำงานในประเทศไทยก็ตาม



รูปที่ 6: การวิเคราะห์ค่าสำคัญจากองค์ประกอบด้านสิ่งอำนวยความสะดวก

จากรูปที่ 6 สิ่งอำนวยความสะดวกภายในอุทยานมีความเหมาะสมกับความต้องการของนักท่องเที่ยว โดยเฉพาะที่พักท่ามกลางธรรมชาติ ห้องพักระยะยาวและมีสิ่งอำนวยความสะดวกพื้นฐานครบครัน และการบริการอาหารมีความหลากหลายในราคาที่เหมาะสม อย่างไรก็ตาม มีบางข้อเสนอแนะที่สะท้อนถึงระยะทางที่ค่อนข้างไกลระหว่างที่พักกับร้านอาหาร หรือมีการใช้เสียงดังในยามวิกาล อาจรบกวนผู้เข้าพักในอุทยาน โดยการบริการของเจ้าหน้าที่ได้รับคำชมในด้านการดูแลรักษาพื้นที่และสัตว์ป่า แต่อาจมีบ้างที่พบกับอุปสรรคทางภาษา นอกจากนี้ยังมีข้อวิจารณ์ถึงการให้ข้อมูลที่ไม่ชัดเจนเกี่ยวกับแผนที่และเส้นทางในอุทยาน หรือขาดข้อมูลที่จำเป็นประโยชน์เกี่ยวกับอันตรายจากสัตว์มีพิษต่าง ๆ ที่อาจพบระหว่างการเดินป่า อย่างไรก็ตาม ทัศนียภาพที่สวยงาม พื้นฐานอย่างห้องน้ำสาธารณะ ไฟฟ้า น้ำประปา สัญญาณอินเทอร์เน็ต หรือจุดชาร์จโทรศัพท์ก็ได้รับการชื่นชมจากนักท่องเที่ยวเช่นกัน



รูปที่ 7: การวิเคราะห์ค่าสำคัญจากองค์ประกอบด้านกิจกรรม

จากรูปที่ 7 กิจกรรมต่าง ๆ ในอุทยานแห่งชาติเขาใหญ่สามารถตอบสนองความต้องการจากนักท่องเที่ยวส่วนใหญ่ได้อย่างน่าพึงพอใจ ด้วยความหลากหลายและเหมาะสมกับสมกับนักท่องเที่ยวในหลากหลายบริบท โดยการเดินป่าเป็นกิจกรรมที่ได้รับความนิยมอย่างมากจากนักท่องเที่ยวชาวต่างชาติ ซึ่งสามารถสัมผัสธรรมชาติและความสะดวกสบายทางชีวภาพได้อย่างใกล้ชิด เช่นเดียวกับการจักรยานที่ชัดเจน แต่บางเส้นทางยังขาดป้ายบอกทางที่ชัดเจน เช่นเดียวกับการส่องสัตว์เป็นอีกกิจกรรมที่น่าสนใจ โดยสามารถพบเห็นสัตว์ป่าร่วมกับการชมธรรมชาติ แต่การมีนักท่องเที่ยวจำนวนมากในช่วงวันหยุดยาวอาจส่งผลกระทบต่อความสะดวกในการใช้บริการ รวมถึงลานกางเต็นท์ในพื้นที่ที่ขาดความเป็นระเบียบ และเสียงรบกวนโดยสัตว์ป่าที่ออกหาอาหารใน

ยามค่าคืน นอกจากนั้น การดูนกในอุทยานที่ได้รับความนิยมอย่างสูงในปัจจุบัน เนื่องจากความหลากหลายของสายพันธุ์นกที่สามารถพบเห็นได้ในพื้นที่ ซึ่งเป็นจุดดึงดูดนักท่องเที่ยวจากทั่วโลก ดังนั้นหากมีไกด์ท้องถิ่นที่มีความรู้จะช่วยให้นักท่องเที่ยวได้รับประสบการณ์ที่ดีขึ้น

5) สรุปผล

พฤติกรรมนักท่องเที่ยวในยุคดิจิทัลมักค้นหาข้อมูลผ่านสื่อออนไลน์ต่าง ๆ เพื่อสนับสนุนการท่องเที่ยว โดยหนึ่งในข้อมูลสำคัญคือคะแนนความพึงพอใจและบทวิจารณ์สาธารณะจากแพลตฟอร์มดิจิทัล แม้ว่าบทวิจารณ์ออนไลน์จะมีข้อดีหลายประการแต่ปริมาณของบทวิจารณ์จำนวนมากทั้งที่เกี่ยวข้องและไม่เกี่ยวข้อง อาจสิ้นเปลืองเวลาและทรัพยากรในการแยกข้อมูลที่สำคัญ ดังนั้นการวิเคราะห์ข้อความอย่างเป็นระบบจึงเป็นวิธีที่ช่วยให้สามารถเข้าใจมุมมองของนักท่องเที่ยวได้อย่างรอบด้าน

งานวิจัยนี้มุ่งเน้นการวิเคราะห์บทวิจารณ์จากนักท่องเที่ยวชาวต่างชาติที่มีต่ออุทยานแห่งชาติเขาใหญ่ด้วยเครื่องมือวิเคราะห์ความรู้สึกแบบใช้พจนานุกรมเป็นฐาน โดยรวบรวมข้อมูลจากบทวิจารณ์สาธารณะเกี่ยวกับอุทยานแห่งชาติเขาใหญ่จากแพลตฟอร์มออนไลน์ระหว่างวันที่ 1 มกราคม พ.ศ. 2563 ถึง 31 ธันวาคม พ.ศ. 2567 รวมทั้งสิ้น 12,035 รายการ เพื่อจำแนกความคิดเห็นของนักท่องเที่ยวออกเป็นเชิงบวก เชิงลบ และเป็นกลางจาก TextBlob Flair และ VADER

ผลการทดสอบ พบว่า VADER มีค่าความถูกต้องเฉลี่ยสูงที่สุดเท่ากับร้อยละ 76 ในการจำแนกความรู้สึกของบทวิจารณ์ อย่างไรก็ตาม ยังพบแนวโน้มของอคติบางประการในกระบวนการวิเคราะห์ เช่น การให้ค่าน้ำหนักกับข้อความเชิงบวกมากเกินไป ส่งผลให้การจำแนกความรู้สึกในบางกรณีไม่ตรงกับความรู้สึกจากบทวิจารณ์ นอกจากนี้ยังพบตัวอย่างของบทวิจารณ์ที่คะแนนความพึงพอใจขัดแย้งกับเนื้อหา เช่น บทวิจารณ์ที่ให้คะแนนสูงแต่มีเนื้อหาวิพากษ์วิจารณ์เกี่ยวกับค่าธรรมเนียมและการให้บริการหรือบทวิจารณ์ที่ให้คะแนนต่ำแต่มีเนื้อหาชื่นชมธรรมชาติของอุทยาน ซึ่งแสดงให้เห็นถึงข้อจำกัดของการใช้เพียงผลการวิเคราะห์ความรู้สึกจากข้อความในการสะท้อนความคิดเห็นของนักท่องเที่ยว ซึ่งอาจไม่เพียงพอ เนื่องจากความคิดเห็นของนักท่องเที่ยวมีความซับซ้อนและหลากหลาย อีกทั้งความคาดหวังของนักท่องเที่ยวแต่ละคนนั้นไม่มีขีดจำกัด และยากที่จะเปรียบเทียบความแตกต่างกันในหลากหลายมุมมองด้วยหลักเกณฑ์ที่เป็นมาตรฐานเดียวกัน

เพื่อลดข้อจำกัดของการใช้ผลการวิเคราะห์ความรู้สึกจากข้อความเพียงอย่างเดียวในการสะท้อนความคิดเห็นของนักท่องเที่ยว และเพิ่มการนำผลลัพธ์ไปใช้ประโยชน์อย่างมีประสิทธิภาพ งานวิจัยนี้วิเคราะห์ข้อมูลเชิงคุณภาพจากบทวิจารณ์ โดยการแปลงข้อความเป็นเวกเตอร์ด้วยวิธี TF-IDF และเปรียบเทียบความสัมพันธ์ระหว่างคำสำคัญและบทวิจารณ์ที่เกี่ยวข้องด้วยการวัดความคล้ายคลึงโคไซน์ ก่อนจะสรุปข้อความด้วยการประมวลผลภาษาธรรมชาติร่วมกับการนำเสนอภาพข้อมูลตามองค์ประกอบด้านการท่องเที่ยว ผลการวิเคราะห์พบว่าแม้ว่าอุทยานแห่งชาติเขาใหญ่จะได้รับการชื่นชมในด้านความงามทางธรรมชาติ แต่ก็ยังมีข้อจำกัดบางประการที่ถูกกล่าวถึงในบทวิจารณ์ของนักท่องเที่ยวชาวต่างชาติ ปัญหาหลักที่พบ ได้แก่ ค่าธรรมเนียมเข้าชมที่แตกต่างกันระหว่างนักท่องเที่ยวชาวไทยและชาวต่างชาติ นอกจากนี้ ระบบขนส่งสาธารณะยังขาดความสะดวกและไม่เพียงพอ ทำให้การเดินทางภายในอุทยานต้องพึ่งพารถส่วนตัวหรือรถเช่าเป็นหลัก ตลอดจนเวลาทำการของอุทยานและข้อจำกัดในการเข้าถึงบางพื้นที่อาจไม่สอดคล้องกับการชมสัตว์ป่าในช่วงเวลาที่แตกต่างกัน

ผลลัพธ์จากงานวิจัยนี้แสดงให้เห็นว่าการประมวลผลภาษาธรรมชาติร่วมกับการวิเคราะห์ความรู้สึกจากข้อความสามารถช่วยให้เข้าใจความคิดเห็นของนักท่องเที่ยวได้ดียิ่งขึ้น ทั้งในแง่ของข้อดีและข้อจำกัดของแหล่งท่องเที่ยว ซึ่งจะนำไปสู่ประโยชน์ต่อการตัดสินใจของนักท่องเที่ยว อีกทั้งยังสามารถนำมากำหนดแนวทางการพัฒนาอุทยานแห่งชาติเขาใหญ่ให้สามารถปรับปรุงคุณภาพการให้บริการและตอบสนองความต้องการของนักท่องเที่ยว นอกจากนี้ ยังเป็นต้นแบบในการประยุกต์ใช้กับแหล่งท่องเที่ยวอื่น ๆ เพื่อยกระดับคุณภาพการให้บริการและเสริมสร้างศักยภาพในการแข่งขันของอุตสาหกรรมการท่องเที่ยวของประเทศไทยในระยะยาว

6) ข้อเสนอแนะ

การวิเคราะห์ความรู้สึกจากข้อความในครั้งต่อไปควรใช้การเรียนรู้เชิงลึกอย่างโมเดล Transformer เพื่อเพิ่มประสิทธิภาพในการประมวลผลข้อมูล นอกจากนี้ ควรมีการประเมินผลจากผู้ใช้งานเพื่อพิจารณาว่าผลลัพธ์จากการวิเคราะห์ข้อความมีความถูกต้องน่าเชื่อถือ และตรงตามความต้องการของกลุ่มเป้าหมาย ทั้งนี้การประเมินจะช่วยให้สามารถปรับปรุงและเพิ่มประสิทธิภาพของผลลัพธ์ก่อนนำไปใช้งานต่อไป

REFERENCES

- [1] Z. Zhu, Y. Zhao, and J. Wang, "The impact of destination online review content characteristics on travel intention: experiments based on psychological distance perspectives," *Aslib J. Inf. Manage.*, vol. 76, no. 1, pp. 42–64, 2024.
- [2] A. Sihombing, L.-W. Liu, and P. Pahrudin, "The impact of online marketing on tourists' visit intention: Mediating roles of trust," *J. Tourism, Heritage Serv. Marketing*, vol. 10, no. 2, pp. 15–23, 2024.
- [3] L. Lei and D. Liu, *Conducting Sentiment Analysis*. Cambridge, U.K.: Cambridge University Press, 2021.
- [4] R. Rachmiani, N. Kintan Oktadinna, and T. Rachmat Fauzan, "The impact of online reviews and ratings on consumer purchasing decisions on e-commerce platforms," *Int. J. Manage. Sci. Inf. Technol.*, vol. 4, no. 2, pp. 504–515, 2024.
- [5] National Statistical Office of Thailand, "Number of tourists in national parks, fiscal year 2022," 2022. [Online]. Available: <https://catalog.dnp.go.th/dataset/91d66c6a-5b88-41ee-8597-11be8aec5aa6/resource/525f31b3-d572-41fa-a822-04f5f3c12ab7/download/tourism65.xlsx>
- [6] National Statistical Office of Thailand, "Number of tourists in national parks, fiscal year 2023," 2023. [Online]. Available: <https://catalog.dnp.go.th/dataset/91d66c6a-5b88-41ee-8597-11be8aec5aa6/resource/70ae333d-19fc-4cdb-a60c-dc43bbf74ad6/download/tourism66.xlsx>
- [7] D. Jurafsky and J. H. Martin, *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*, 2nd ed. Upper Saddle River, NJ, USA: Prentice Hall, 2008.
- [8] T. Snake, *Natural Language Processing Practical Using Transformers with Python*. 2022.
- [9] L. T. Dela Cruz *et al.*, "An enhancement of support vector machine in the context of sentiment analysis applied in scraped data from Tripadvisor hotel reviews," *Int. J. Comput. Sci. Res.*, vol. 8, pp. 3235–3251, 2024.
- [10] S. Jardim and C. Mora, "Customer reviews sentiment-based analysis and clustering for market-oriented tourism services and products development or positioning," *Procedia Comput. Sci.*, vol. 196, pp. 199–206, 2022.
- [11] K. Puh and M. Bagić Babac, "Predicting sentiment and rating of tourist reviews using machine learning," *J. Hosp. Tour. Insights*, vol. 6, no. 3, pp. 1188–1204, 2023.
- [12] I. Perikos and A. Diamantopoulos, "Explainable aspect-based sentiment analysis using transformer models," *Big Data Cogn. Comput.*, vol. 8, no. 11, 2024, Art. no. 141.
- [13] K. W. Trisna, J. Huang, H. Lei, and E. M. Dharma, "From context-independent embedding to transformer: Exploring sentiment classification in online reviews with deep learning approaches," *J. Theor. Appl. Inf. Technol.*, vol. 102, no. 19, pp. 6980–7003, 2024.
- [14] P. Suanpang, P. Jamjuntr, and P. Kaewyong, "Sentiment analysis with a TextBlob package Implications for tourism," *J. Manage. Inf. Decis. Sci.*, vol. 24, no. S6, pp. 1–9, 2021.
- [15] I. Khumtaveeporn and K. Wattanasuwan, "AI sentiment analysis for destination branding: A case study of Buriram, Thailand," *J. Bus. Adm.*, vol. 46, no. 180, pp. 50–74, 2023.
- [16] K. Barik and S. Misra, "Analysis of customer reviews with an improved VADER lexicon classifier," *J. Big Data*, vol. 11, 2024, Art. no. 10.
- [17] D. Buhalis, "Marketing the competitive destination of the future," *Tourism Manage.*, vol. 21, no. 1, pp. 97–116, 2000.
- [18] R. Mitchell, *Web Scraping with Python*, 3rd ed. Sebastopol, CA, USA: O'Reilly Media, 2024.
- [19] M. Arief and N. A. Samsudin, "Hybrid approach with VADER and multinomial logistic regression for multiclass sentiment analysis in online customer review," *Int. J. Adv. Comput. Sci. Appl.*, vol. 14, no. 12, pp. 311–320, 2023.
- [20] D. Tay, *Data Analytics for Discourse Analysis with Python: The Case of Therapy Talk*. New York, NY, USA: Routledge, 2024.
- [21] B. Liu, *Sentiment Analysis and Opinion Mining*. Cham, Switzerland: Springer, 2022.

การประยุกต์ใช้โครงข่ายทรานส์ฟอร์มเมอร์สำหรับจำแนกและจดจำโรคเล็บ

เอกรัตน์ สุขสุคนธ์^{1*} บุญธิดา ชุนงาม² เอกชัย เนาวนิช³

^{1,2}สาขาวิชาวิศวกรรมคอมพิวเตอร์ คณะครุศาสตร์อุตสาหกรรม มหาวิทยาลัยเทคโนโลยีราชมงคลสุวรรณภูมิ, สุพรรณบุรี, ประเทศไทย

³สาขาวิชาเทคโนโลยีดิจิทัลมีเดีย คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยเทคโนโลยีราชมงคลสุวรรณภูมิ, นนทบุรี, ประเทศไทย

*ผู้ประพันธ์บรรณกิจ อีเมล : aekkarat.s@rmutsb.ac.th

รับต้นฉบับ : 13 กุมภาพันธ์ 2568; รับบทความฉบับแก้ไข : 6 มีนาคม 2568; ตอบรับบทความ : 8 สิงหาคม 2568

เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

โรคเล็บเป็นปัญหาสำคัญในการวินิจฉัยทางการแพทย์ เนื่องจากโรคเล็บมีลักษณะที่คล้ายคลึงกัน ทำให้การวินิจฉัยต้องอาศัยผู้เชี่ยวชาญด้านโรคผิวหนัง ความผิดพลาดในการวินิจฉัยอาจนำไปสู่การรักษาที่ผิดพลาดและส่งผลกระทบต่อผู้ป่วยมีอาการแทรกซ้อนเพิ่มขึ้น งานวิจัยนี้นำเสนอการประยุกต์ใช้โครงข่ายประสาทเทียมแบบทรานส์ฟอร์มเมอร์สำหรับจำแนกโรคเล็บ เนื่องจากโครงข่ายมีความสามารถในการเรียนรู้คุณลักษณะที่สำคัญซึ่งก่อให้เกิดโรค แนวทางนี้เป็นการส่งเสริมเพื่อเพิ่มความแม่นยำในการจำแนก การทดลองมุ่งเน้นศึกษาตัวอย่างโรคเล็บที่หลากหลาย เช่น เล็บที่เป็นโรคสะเก็ดเงิน เล็บติดเชื้อรา และเล็บปกติ โดยโครงข่ายได้รับการฝึกอบรมด้วยการปรับแต่งค่าพารามิเตอร์การเรียนรู้เพื่อเพิ่มประสิทธิภาพการทำงาน โดยประเมินด้วยค่าความแม่นยำมีผลสูงสุดจากการฝึกอบรมที่ร้อยละ 99.40 ซึ่งแสดงถึงศักยภาพในการจำแนกและจดจำโรคเล็บได้อย่างมีประสิทธิภาพ แนวทางนี้ไม่เพียงแต่ช่วยเพิ่มความแม่นยำในการวินิจฉัย แต่ยังสามารถพัฒนาระบบที่ลดภาระงานของบุคลากรทางการแพทย์ และช่วยให้การวินิจฉัยโรคเป็นไปได้อย่างรวดเร็ว ซึ่งอาจนำไปสู่การพัฒนากระบวนการวินิจฉัยอัตโนมัติที่ช่วยยกระดับคุณภาพการดูแลและรักษาในอนาคต

คำสำคัญ : ปัญญาประดิษฐ์ การเรียนรู้เชิงลึก การจำแนกและจดจำโรคเล็บ โครงข่ายทรานส์ฟอร์มเมอร์

Exploiting Transformer Network for Nail Diseases Classification and Recognition

Aekkarat Suksukont^{1*} Bunthida Chunngam² Ekachai Naowanich³

^{1,2}*Department of Computer Engineering, Faculty of Industrial Education,
Rajamangala University of Technology Suvarnabhumi, Suphan Buri, Thailand*

³*Department of Digital Media Technology, Faculty of Science and Technology,
Rajamangala University of Technology Suvarnabhumi, Nonthaburi, Thailand*

*Corresponding Author. E-mail address: aekkarat.s@rmutsb.ac.th

Received: 13 February 2025; Revised: 6 March 2025; Accepted: 8 August 2025

Published online: 26 December 2025

Abstract

Diagnosing nail diseases is a complex task due to their similar visual characteristics, often requiring expert dermatologists for accurate assessment. Misdiagnosis can lead to ineffective treatment and prolonged patient discomfort. This study explores the use of a transformer neural network for classifying nail diseases, leveraging its ability to identify intricate patterns and subtle features that may indicate early signs of disease. The research focuses on three nail conditions: psoriasis nails, onychomycosis, and healthy. The model was trained with a carefully optimized set of hyperparameters to improve learning efficiency and classification performance. Experimental results showed that the network achieved a peak accuracy of 99.40%, demonstrating its ability to effectively distinguish between different nail conditions. This approach not only enhances classification accuracy but also has the potential to reduce the workload of healthcare professionals and speed up diagnosis. Ultimately, this advancement could contribute to the development of automated diagnostic systems, leading to improved patient care and treatment outcomes.

Keywords: Artificial intelligence, Deep learning, Nail disease classification and recognition, Transformer network

1) บทนำ

โรคเล็บเป็นปัญหาสุขภาพที่พบได้ทั่วไปและอาจส่งผลต่อคุณภาพชีวิตของผู้ป่วย โดยมีสาเหตุจากหลายปัจจัย เช่น การติดเชื้อรา การอักเสบจากแบคทีเรีย โรคผิวหนัง หรือภาวะผิดปกติของร่างกาย แม้ว่าการวินิจฉัยโรคเล็บจะต้องอาศัยความเชี่ยวชาญของแพทย์ผิวหนัง แต่เนื่องจากลักษณะของโรคมีความคล้ายคลึงกันและซับซ้อน การตรวจวินิจฉัยแบบดั้งเดิมจึงมักใช้เวลานานและมีข้อจำกัดในการเข้าถึงผู้เชี่ยวชาญ [1], [2] ในปัจจุบันเทคโนโลยีปัญญาประดิษฐ์ (artificial intelligence) และการเรียนรู้เชิงลึก (deep learning) ได้เข้ามามีบทบาทสำคัญในการแพทย์ โดยเฉพาะการนำเทคนิคการมองเห็นของคอมพิวเตอร์มาประยุกต์ใช้ในการวิเคราะห์ภาพทางการแพทย์เพื่อช่วยในการตรวจวินิจฉัยโรค การนำเทคโนโลยีเหล่านี้มาช่วยจำแนกและจดจำโรคเล็บสามารถเพิ่มความแม่นยำ ลดข้อผิดพลาด และช่วยให้แพทย์สามารถให้การวินิจฉัยได้รวดเร็วขึ้น [3], [4]

ด้วยเหตุนี้ผู้วิจัยจึงมุ่งเน้นการพัฒนาแบบจำลองการเรียนรู้เชิงลึกที่สามารถจำแนกประเภทของโรคเล็บโดยอัตโนมัติ เพื่อเพิ่มความแม่นยำและประสิทธิภาพของการวินิจฉัย ซึ่งอาจนำไปสู่การลดภาระงานของแพทย์และเพิ่มโอกาสในการเข้าถึงการรักษาที่รวดเร็วและมีประสิทธิภาพยิ่งขึ้น

2) ทบทวนวรรณกรรม

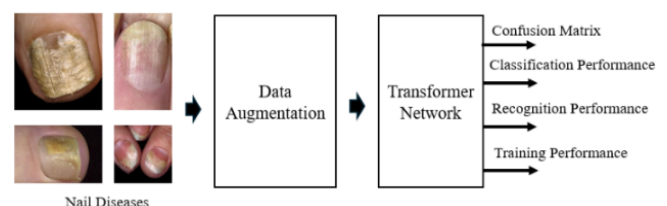
ในช่วงไม่กี่ปีที่ผ่านมา เทคนิคการเรียนรู้เชิงลึกโดยเฉพาะโครงข่าย convolutional neural network (CNN) ที่มีบทบาทสำคัญในการประยุกต์ใช้สำหรับการวิเคราะห์ภาพถ่ายทางการแพทย์ [5] เนื่องจากความสามารถในการเรียนรู้ และการตรวจจับคุณลักษณะเชิงลึกได้อย่างมีประสิทธิภาพ [6] และยังได้รับการพัฒนาให้สามารถเรียนรู้องค์ประกอบจากส่วนย่อย ๆ จากความสัมพันธ์เชิงพื้นที่ของเป้าหมายนั้น ๆ ส่งผลให้มีการพัฒนาเพื่อแก้ปัญหาภาพถ่ายที่ซับซ้อน เช่น ภาพถ่ายรังสี [7] ภาพสแกน MRI [8] และภาพทางผิวหนัง [9] โดยเฉพาะการนำ CNN มาประยุกต์ใช้ในการจำแนกโรคเล็บเพื่อช่วยเพิ่มความแม่นยำและลดข้อผิดพลาดในการวินิจฉัย เช่น ใน Hariyani et al. [4] เสนอการแบ่งส่วนเส้นเลือดฝอยที่เล็บ มุ่งเน้นสัญญาณรบกวนและความแปรปรวนสูง โดยอาศัย manual, semi-automated และ automated ในการแบ่งส่วนภาพ ร่วมกับ DA-CapNet บน U-Net ที่ปรับปรุงด้วย dual attention module เพื่อช่วยให้โมเดลสามารถเรียนรู้คุณลักษณะของภาพได้ดีขึ้น ใน Yamaç et al. [10]

เสนอ ResNet101v2 ร่วมกับการถ่ายโอนการเรียนรู้เพื่อจำแนกลักษณะของโรคเล็บที่หลากหลาย เพื่อศึกษาการเรียนรู้ด้วยข้อมูลที่จำกัดและคุณภาพต่ำ ใน Muneera Begum et al. [11] เสนอ CNN โดยพัฒนาให้อยู่ในรูปของเว็บแอปพลิเคชัน DermaDoc เพื่อช่วยผู้ใช้ตรวจสอบโรคและให้ข้อมูลเกี่ยวกับอาการและแนวทางการรักษาเบื้องต้น และ Shandilya et al. [12] แสดงให้เห็นถึงความสำคัญของการประยุกต์ใช้และพัฒนา อย่างไรก็ตาม การพัฒนายังคงมีข้อด้อยที่เกิดจากการทำงานของโครงข่ายสัญญาณรบกวนภาพ และความซับซ้อนของโรคเล็บ ซึ่งชี้ให้เห็นแนวทางในการพัฒนาเพื่อเพิ่มประสิทธิภาพ

งานวิจัยนี้นำเสนอการประยุกต์ใช้โครงข่ายทรานส์ฟอร์มเมอร์สำหรับจำแนกโรคเล็บ โดยโมเดลได้รับการฝึกอบรมด้วยชุดข้อมูลภาพเล็บที่ผ่านการตรวจสอบจากแพทย์ผู้เชี่ยวชาญ พร้อมทั้งดำเนินการปรับปรุงการเรียนรู้ของโครงข่ายในขณะฝึกอบรม ประสิทธิภาพของโมเดลถูกประเมินโดยค่า accuracy, precision, recall, F1 และ confusion matrix เทคนิคและวิธีที่นำเสนอเป็นการพัฒนาเพื่อเพิ่มความสามารถของโครงข่าย ให้สามารถนำไปพัฒนาควบคู่กับระบบการแพทย์ทางไกล เพื่อลดภาระงานของแพทย์และเพิ่มการเข้าถึงในการรักษา

3) วิธีดำเนินการวิจัย

งานวิจัยนี้ดำเนินการโดยรวบรวมชุดข้อมูลภาพเล็บที่ผ่านการตรวจสอบจากชุดข้อมูลออนไลน์ที่เผยแพร่ไว้สำหรับนำไปพัฒนาปัญญาประดิษฐ์ โดยนำมาปรับแต่งภาพด้วยการเสริมข้อมูล (data augmentation) เพื่อเพิ่มความหลากหลายของข้อมูล จากนั้นนำมาฝึกอบรมโครงข่ายทรานส์ฟอร์มเมอร์ โดยทำปรับเปลี่ยนพารามิเตอร์ (fine-tuning) บนชุดข้อมูลโรคเล็บ การฝึกอบรมโครงข่ายได้รับการประเมินด้วย accuracy, precision, recall, F1 และ confusion matrix พร้อมเปรียบเทียบกับโมเดลก่อนหน้า การทำงานดังกล่าวแสดงไว้ดังรูปที่ 1



รูปที่ 1 : ขั้นตอนการจำแนกและจดจำโรคเล็บ

3.1) ชุดข้อมูล (Dataset)

Nail disease [13] เป็นชุดข้อมูลที่เผยแพร่สาธารณะ ถูกพัฒนาขึ้นสำหรับการศึกษาวิจัยในการจำแนกและจดจำโรคเล็บ ประกอบด้วยสองชุดย่อย ได้แก่ ชุดเทรนนิ่ง (training set) ร้อยละ 80 และชุดทดสอบร้อยละ 20 โดยข้อมูลแต่ละชุดแบ่งออกเป็นสามประเภท ได้แก่ เล็บปกติ (healthy), โรคเชื้อรา (onychomycosis) และ โรคสะเก็ดเงิน (psoriasis) ชุดข้อมูลนี้ได้รับการจัดระเบียบเพื่อใช้ในการฝึกและทดสอบโมเดลของการเรียนรู้เชิงลึก โดยสามารถดาวน์โหลดได้จาก Kaggle เพื่อนำไปใช้ในการพัฒนาระบบวินิจฉัยโรคเล็บแบบอัตโนมัติ



รูปที่ 2 : ชุดข้อมูลโรคเล็บ

3.2) โครงข่ายทรานส์ฟอร์มเมอร์ (Transformer Network)

โครงข่ายทรานส์ฟอร์มเมอร์ [14] เป็นสถาปัตยกรรมที่ใช้ self-attention mechanism ถูกออกแบบเพื่อประมวลผลข้อมูลที่เป็นลำดับส่งผลให้โมเดลสามารถตรวจจับลักษณะเด่นของข้อมูลภาพได้โดยไม่ต้องอาศัยโครงสร้างแบบลำดับเชิงพื้นที่ หรือความสัมพันธ์ระหว่างส่วนต่างๆ ของภาพ มีการทำงานโดยใช้โครงข่ายที่แบ่งภาพออกเป็นแพตช์ ขนาด $P \times P$ และแปลงเป็นเวกเตอร์ผ่าน linear projection จากนั้น positional encoding จะถูกเพิ่มเข้าไปเพื่อรักษาข้อมูลลำดับของแพตช์ในกระบวนการ self-attention ภายใน multi-head self-attention (MHSA) ซึ่งทำหน้าที่รวมองค์ประกอบของ scaled dot-product attention เพื่อเรียนรู้ความสัมพันธ์เชิงลึกของข้อมูล ดังสมการที่ (1)

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (1)$$

เมื่อ K , Q , V คือ เวกเตอร์ query, key, value ที่ได้จากการแปลงข้อมูลผ่านเมทริกซ์น้ำหนัก และ d_k คือ มิติของ key เพื่อปรับสเกลให้เหมาะสม MHSA ซึ่งจะดำเนินการไปพร้อมกันด้วย self-attention เพื่อนำข้อมูลเชิงลึกไปใช้หลังจากผ่านขั้นตอนของ feedforward network และ layer normalization โดยโมเดลจะรวมคุณลักษณะและส่งต่อไปยัง โครงข่ายประสาทเพื่อจำแนกประเภทและจดจำ

3.3) การวัดประสิทธิภาพ

ประสิทธิภาพของโมเดลสามารถใช้การคำนวณด้วย accuracy เพื่อประเมินภาพรวมของการฝึกอบรม ดังสมการที่ (2), precision เพื่อคำนวณจากสัดส่วนของ true positive (TP) ต่อผลรวมของ TP และ false positive (FP) ซึ่งใช้วัดความแม่นยำของโมเดลในการทำนายผลลัพธ์ที่เป็นบวก ดังสมการที่ (3), recall เป็นการคำนวณจากสัดส่วนของ TP ต่อผลรวมของ TP และ false negative (FN) ซึ่งใช้วัดความสามารถของโมเดลในการดึงข้อมูลที่ถูกต้อง ดังสมการที่ (4), F1 ผลทางสถิติที่ใช้วัดความแม่นยำของโมเดล ซึ่งเป็นค่าเฉลี่ยน้ำหนักระหว่าง precision และ recall ดังสมการที่ (5) ในขณะที่ confusion matrix ถูกนำมาอธิบายชุดข้อมูลที่ไม่สมดุล

$$\text{Accuracy} = \frac{TP + TN}{TP + FP + FN + TN} \quad (2)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (3)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (4)$$

$$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (5)$$

โดยที่ TP และ true negative (TN) คือ จำนวนตัวอย่างที่จำแนกและจดจำได้ถูกต้อง FP และ FN คือ จำนวนตัวอย่างที่จำแนกและจดจำผิดพลาด

3.4) อุปกรณ์ในการทดลอง

คอมพิวเตอร์ที่ใช้ในการทดลองดำเนินการบนระบบปฏิบัติการ Windows ที่ติดตั้ง CPU Intel Core i5-12400F LGA 1700 @ 2.5 GHz พร้อมด้วย RAM ขนาด 32 GB ที่มีความเร็วสัญญาณนาฬิกา 5600 MHz และ GPU NVIDIA RTX 4070 ซึ่งมาพร้อม VRAM ขนาด 12 GB และ CUDA cores จำนวน 5888 หน่วย

NumPy และ TensorFlow เป็นเครื่องมือในการพัฒนาและทดสอบ การกำหนดค่าพารามิเตอร์สำหรับการเทรนนิ่ง และทดสอบเพื่อช่วยเพิ่มความเร็วในการประมวลผล ถูกกำหนดไว้เพื่อความเป็นมาตรฐาน ดังสรุปในตารางที่ 1

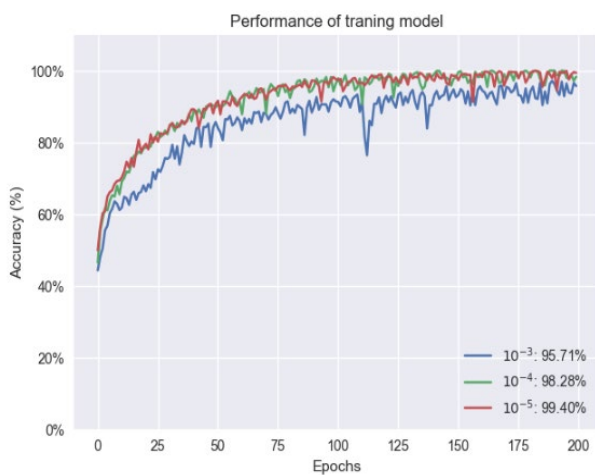
ตารางที่ 1 : พารามิเตอร์สำหรับการเทรนนิ่งโมเดล

พารามิเตอร์	ค่าของพารามิเตอร์
ขนาดข้อมูล	150×150×3
อัตราการเรียนรู้	10^{-3} , 10^{-4} , 10^{-5}
รอบของการเทรนนิ่ง	200
การเพิ่มประสิทธิภาพ	Adam

4) ผลการวิจัยและอภิปรายผล

4.1) ผลการเทรนนิ่ง (Training Results)

ผลของการเทรนนิ่ง ในรูปที่ 3 แสดงประสิทธิภาพการเรียนรู้ของโครงข่ายในขณะที่มีการเปลี่ยนแปลงอัตราการเรียนรู้ 10^{-3} , 10^{-4} และ 10^{-5} ซึ่งผลการเทรนนิ่ง พบว่า 10^{-5} ให้ประสิทธิภาพการเทรนนิ่ง สูงสุดที่ร้อยละ 99.40 และมีความเสถียรมากที่สุดในขณะที่ 10^{-4} มีความสมดุลในการเรียนรู้โดยมีผลสูงสุดที่ร้อยละ 98.28 ซึ่งเป็นผลเทียบเท่า 10^{-5} และ 10^{-3} มีผลการเรียนรู้ที่ผันผวนโดยมีผลสูงสุดที่ร้อยละ 95.71



รูปที่ 3 : ประสิทธิภาพการเทรนนิ่งโมเดล

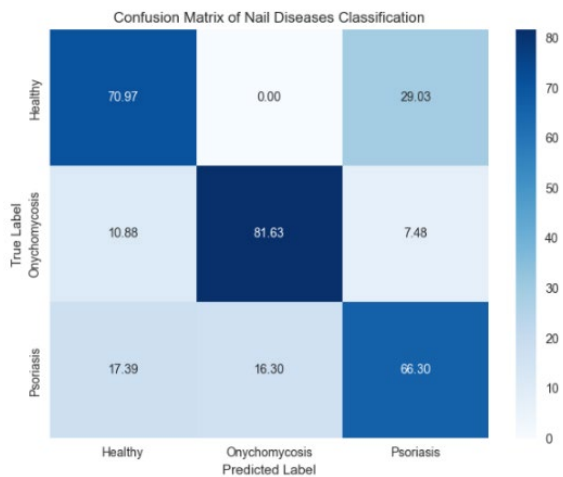
4.2) ผลการจำแนกและจดจำ

เมื่อนำโครงข่ายไปทดสอบ โดยนำ 10^{-5} ที่ได้รับการเทรนนิ่ง สูงสุดไปทดสอบ โดยใช้ตัวชี้วัดค่า precision, recall และ F1 ผลการจำแนกโรคเล็บ ได้แก่ เล็บปกติ, โรคเชื้อรา, และ โรคสะเก็ดเงิน สรุปผลในตารางที่ 2 ผลของโรคเชื้อราได้รับผลสูงสุด

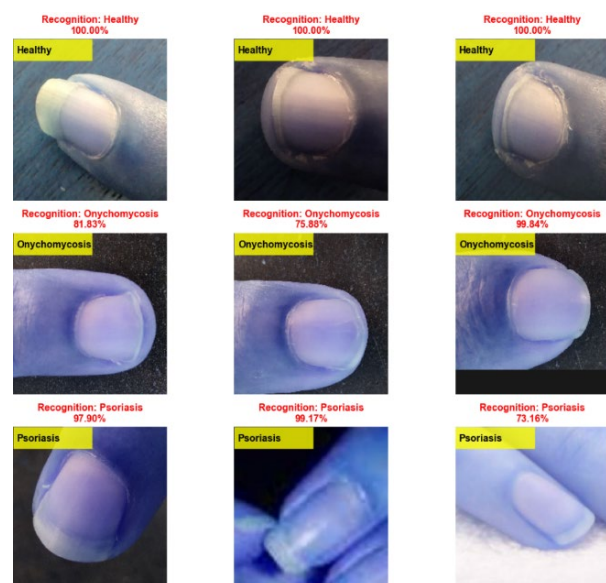
ที่ร้อยละ 88.90, 81.60 และ 85.10 ของ precision, recall และ F1 สะท้อนถึงความสมดุลในการจำแนก อย่างไรก็ตาม การจำแนกเล็บปกติและโรคสะเก็ดเงินมีผลค่อนข้างน้อยเมื่อเทียบกับการทดลองชุดอื่น ๆ โดยสามารถจำแนกได้เพียงร้อยละ 57.90, 71.00 และ 63.80 ของเล็บปกติ และมีผลที่ร้อยละ 67.80, 66.30 และ 67.00 ของโรคสะเก็ดเงิน ตามลำดับ

ตารางที่ 2 : ประสิทธิภาพการจำแนกโรคเล็บ

โรคเล็บ	Precision (%)	Recall (%)	F1 (%)
เล็บปกติ	57.90	71.00	63.80
โรคเชื้อรา	88.90	81.60	85.10
โรคสะเก็ดเงิน	67.80	66.30	67.00



รูปที่ 4 : Confusion matrix ของการจำแนกโรคเล็บ



รูปที่ 5: ผลการจดจำโรคเล็บ

ในรูปที่ 4 ผลลัพธ์ของ 10^{-5} เมื่อนำมาทดสอบเพื่อจำแนกโรคเล็บ โดย confusion matrix แสดงผลการจำแนกโรคเล็บ ผลสูงสุดที่ร้อยละ 70.97 สำหรับเล็บปกติ, ร้อยละ 81.63 สำหรับโรคเชื้อรา และร้อยละ 66.30 สำหรับโรคสะเก็ดเงิน อย่างไรก็ตาม ในการทดสอบยังมีข้อผิดพลาดที่เกิดจากความสับสนในการเรียนรู้ข้อมูล เช่น ผลการจำแนกที่ร้อยละ 29.03 ของเล็บปกติ จำแนกผิดเป็นโรคสะเก็ดเงิน หรือผลการจำแนกที่ร้อยละ 17.39 ของโรคสะเก็ดเงินถูกทำนายผิดเป็นเล็บปกติ

ในรูปที่ 5 ผลลัพธ์ของ 10^{-5} เมื่อนำมาทดสอบเพื่อจดจำโรคเล็บ สามารถแสดงผลการจดจำได้ทั้ง เล็บปกติ, โรคเชื้อรา และโรคสะเก็ดเงิน โดยมีความแม่นยำสูงในบางภาพ เช่น ร้อยละ 100 สำหรับเล็บปกติ และร้อยละ 99.84 สำหรับโรคเชื้อราที่เล็บ อย่างไรก็ตาม ในบางกรณีมีค่าความแม่นยำต่ำ เช่น ร้อยละ 73.16 สำหรับโรคสะเก็ดเงิน ซึ่งอาจแสดงถึงความผิดพลาดในบางกรณี

4.3) การเปรียบเทียบผลการทดลอง

งานวิจัยนี้มีวัตถุประสงค์เพื่อประยุกต์ใช้โครงข่ายทรานส์ฟอร์มเมอร์สำหรับจำแนกและจดจำโรคเล็บ โดยได้รวบรวมผลการศึกษานี้ เพื่อประเมินประสิทธิภาพของวิธีที่เสนอ สรุปผลในตารางที่ 3

ตารางที่ 3 : การเปรียบเทียบประสิทธิภาพการเทรนนิ่ง (training) กับการศึกษาก่อนหน้า

บทความ	เทคนิค	เทรนนิ่ง (%)
[4]	DA-CapNet	81.00 (recall)
[10]	ResNet101v2	89.00
[12]	Hybrid Capsule CNN	99.40
[15]	Hybrid CNN	80.45
[16]	Hybrid CNN-RF	96.08
[17]	CNN	87.33
งานวิจัยนี้	เทคนิคที่นำเสนอ	99.40

ในตารางที่ 3 เปรียบเทียบประสิทธิภาพกับการศึกษาก่อนหน้า ผลลัพธ์ของเทคนิค DA-CapNet, ResNet101v2, Hybrid Capsule CNN, Hybrid CNN, CNN-RF และ CNN มีผลการฝึกอบรมที่ได้ โดยเทคนิคที่นำเสนอในงานวิจัยนี้ใช้ Transformer ซึ่งให้ผลลัพธ์ที่ดีที่สุดที่ร้อยละ 99.40 แสดงให้เห็นว่าเทคนิคที่นำเสนอมีประสิทธิภาพดีกว่าเมื่อเทียบกับวิธีอื่น ๆ

5) อภิปรายผลการทดลอง

ผลการทดลองแสดงให้เห็นว่าอัตราการเรียนรู้แตกต่างกัน ส่งผลต่อประสิทธิภาพของโครงข่ายที่เสนอ ซึ่งผลของ 10^{-5} ให้ผลลัพธ์ที่ดีที่สุดที่ร้อยละ 99.40 และมีความเสถียรมากที่สุด แสดงให้เห็นว่าอัตราการเรียนรู้ที่น้อยช่วยให้โครงข่ายเรียนรู้คุณลักษณะของข้อมูลได้ดียิ่งขึ้น เมื่อเทียบกับพารามิเตอร์อื่นๆ ทั้งนี้เมื่อสังเกต 10^{-3} มีผลลัพธ์ที่ผันผวนมากที่สุด ซึ่งบ่งชี้ว่าอัตราการเรียนรู้ที่สูงเกินไปทำให้การเรียนรู้เกิดปัญหาการเรียนรู้ที่ไม่มีเสถียรภาพ

เมื่อทดสอบกับชุดข้อมูลเพื่อจำแนกโรคเล็บพบว่า อัตราการเรียนรู้ที่ 10^{-5} สามารถจำแนกโรคเชื้อราได้ดีที่สุด โดยมีค่า precision สูงสุดที่ร้อยละ 88.90 มีผล recall สูงสุดที่ร้อยละ 81.60 และ F1 สูงสุดที่ร้อยละ 85.10 ซึ่งสะท้อนถึงความสามารถในการจำแนกที่เฉพาะในบางกรณี นอกจากนี้ เมื่อทดสอบเพื่อจดจำโรคเล็บ โมเดลสามารถจดจำเล็บปกติได้อย่างแม่นยำที่ร้อยละ 100 ในทุกตัวอย่าง สำหรับโรคเชื้อรามีความแม่นยำแตกต่างกัน ซึ่งมีผลสูงสุดที่ร้อยละ 99.84 และในบางกรณีมีผลลดลงที่ร้อยละ 75.80 ทั้งนี้ ผลการทดลองเป็นไปในแนวทางเดียวกันกับโรคสะเก็ดเงินที่มีความแม่นยำสูงและมีผลลดน้อยลงในบางกรณี อาจเกิดจากความคล้ายคลึงกันของลักษณะทางกายภาพของโรคสะท้อนถึงข้อจำกัดของโมเดลในการจำแนกและจดจำโรคบนสภาพผิวของเล็บ

แม้ว่าผลลัพธ์โดยรวมจะแสดงให้เห็นว่าโครงข่ายที่เสนอมีประสิทธิภาพสูงในการจำแนกและจดจำโรคเล็บ แต่ยังสามารถพัฒนาเพิ่มเติม โดยเฉพาะการเพิ่มชุดข้อมูลที่ครอบคลุมตัวอย่างที่มีลักษณะคล้ายคลึงกัน หรือการปรับแต่งโครงข่ายให้สามารถดึงลักษณะเฉพาะของโรคให้ดียิ่งขึ้น

6) สรุปผลการวิจัย

งานวิจัยนี้นำเสนอการประยุกต์ใช้โครงข่ายทรานส์ฟอร์มเมอร์สำหรับจำแนกและจดจำโรคเล็บ เนื่องจากข้อดีในการเรียนรู้ความสัมพันธ์เชิงพื้นที่ได้ดีและลดการสูญเสียข้อมูลที่ซับซ้อนฐานข้อมูลที่ใช้ทดลองประกอบด้วยโรคเล็บ ได้แก่ เล็บปกติ โรคเชื้อราที่เล็บ และโรคสะเก็ดเงิน โดยโมเดลได้รับการฝึกและปรับแต่งพารามิเตอร์เพื่อให้มีประสิทธิภาพสูงสุดสำหรับการเทรนนิ่งและทดสอบ

ผลการทดลองแสดงให้เห็นว่า วิธีที่เสนอมีประสิทธิภาพการเทรนนิ่งสูงสุดที่ร้อยละ 99.40 อีกทั้งยังสามารถจำแนกและจดจำ

โรคเล็บได้อย่างแม่นยำ ทำให้เหมาะสมสำหรับนำไปพัฒนาระบบวินิจฉัยโรคเล็บแบบอัตโนมัติ

จากผลการทดลองสรุปได้ว่า โครงข่ายทรานส์ฟอร์มเมอร์เป็นแนวทางที่มีประสิทธิภาพสูงในการจำแนกและจดจำโรคเล็บ โดยช่วยปรับปรุงความแม่นยำในการวินิจฉัย ซึ่งสามารถนำไปพัฒนาระบบเพื่อลดภาระงานของแพทย์ในอนาคต

7) ข้อเสนอแนะ

สำหรับการพัฒนาในอนาคต ควรขยายฐานข้อมูลให้ครอบคลุมโรคเล็บที่หลากหลายมากขึ้นและเพิ่มคุณภาพของข้อมูลเพื่อปรับปรุงความแม่นยำของโมเดล นอกจากนี้ควรพัฒนาโครงข่ายทรานส์ฟอร์มเมอร์ให้มีประสิทธิภาพสูงขึ้นโดยใช้เทคนิคใหม่ ๆ เช่น hybrid transformer-CNN หรือ self-supervised learning รวมถึงการพัฒนาแอปพลิเคชันสำหรับการใช้งานจริงผ่านสมาร์ตโฟนและการผสมรวมกับระบบการแพทย์ทางไกล เพื่อให้แพทย์และผู้ป่วยสามารถเข้าถึงการวินิจฉัยได้ดียิ่งขึ้น

REFERENCES

- [1] M. Umei *et al.*, "Prevalence, risk factors and potential implications of nail biting in adults with congenital heart disease," *Int. J. Cardiol.*, vol. 418, 2025, Art. no. 132652, doi: 10.1016/j.ijcard.2024.132652.
- [2] K. L. Curtis *et al.*, "Diagnosis and management of nail unit squamous cell carcinoma: A clinical review by an expert panel," *JAAD Rev.*, vol. 4, pp. 187–194, Jun. 2025, doi: 10.1016/j.jdrv.2024.12.012.
- [3] S. Das and M. Bhattacharjee, "Image-based sensing of Leukonychia for early diagnosis of anemia using a smartphone application," *IEEE Sens. Lett.*, vol. 6, no. 11, 2022, Art. no. 3502604.
- [4] Y. S. Hariyani, H. Eom, and C. Park, "DA-capnet: dual attention deep learning based on U-net for nailfold capillary segmentation," *IEEE Access*, vol. 8, pp. 10543–10553, 2020, doi: 10.1109/ACCESS.2020.2965651.
- [5] A. Prommakhot and J. Srinonchat, "Combining convolutional neural networks for fungi classification," *IEEE Access*, vol. 12, pp. 58021–58030, 2024, doi: 10.1109/ACCESS.2024.3391630.
- [6] A. Prommakhot and J. Srinonchat, "VGGNet integration for kidney tumor classification," in *Proc. 12th Int. Electr. Eng. Congr. (IEECON)*, Pattaya, Thailand, Mar. 2024, pp. 1–6.
- [7] H. Wang, H. Jia, L. Lu, and Y. Xia, "Thorax-net: An attention regularized deep neural network for classification of Thoracic diseases on chest radiography," *IEEE J. Biomed. Health Inform.*, vol. 24, no. 2, pp. 475–485, 2020.
- [8] A. P. Leynes, N. Deveshwar, S. S. Nagarajan, and P. E. Z. Larson, "Scan-specific self-supervised bayesian deep non-linear inversion for undersampled MRI reconstruction," *IEEE Trans. Med. Imaging*, vol. 43, no. 6, pp. 2358–2369, 2024.
- [9] K. Lee, T. C. Cavalcanti, S. Kim, H. M. Lew, D. H. Suh, and D. H. Lee, "Multi-task and few-shot learning-based fully automatic deep learning platform for mobile diagnosis of skin diseases," *IEEE J. Biomed. Health Inform.*, vol. 27, no. 1, pp. 176–187, 2023.
- [10] S. A. Yamaç, O. Kuyucuoglu, Ş. B. Köseoglu, and S. Ulukaya, "Deep learning-based classification of human nail diseases using color nail images," in *Proc. 45th Int. Conf. Telecommun. and Signal Process. (TSP)*, Prague, Czech Republic, Jul. 2022, pp. 196–199.
- [11] H. Muneera Begum, A. Dhivya, A. J. Krishnan, and S. D. Keerthana, "Automated detection of skin and nail disorders using convolutional neural networks," in *Proc. 5th Int. Conf. Trends Electron. Inform. (ICOEI)*, Tirunelveli, India, Jun. 2021, pp. 1309–1316.
- [12] G. Shandilya *et al.*, "Autonomous detection of nail disorders using a hybrid capsule CNN: A novel deep learning approach for early diagnosis," *BMC Med. Inform. Decis. Mak.*, vol. 24, 2024, Art. no. 414.
- [13] *Nail Disease Image Classification Dataset*, Kaggle, 2023. [Online]. Available: <https://www.kaggle.com/datasets/josephrasanjana/nail-disease-image-classification-dataset>
- [14] A. Vaswani *et al.*, "Attention is all you need," in *Proc. 31st Conf. Neural Inf. Process. Syst. (NIPS)*, Long Beach, CA, USA, Dec. 2017, pp. 5998–6008. [Online]. Available: https://papers.nips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf
- [15] R. Nijhawan, R. Verma, Ayushi, S. Bhushan, R. Dua, and A. Mittal, "An integrated deep learning framework approach for nail disease identification," in *Proc. 13th Int. Conf. Signal-Image Technol. Internet-Based Syst. (SITIS)*, Jaipur, India, Dec. 2017, pp. 197–202.
- [16] A. Kumar and K. K. Tiwari, "Disease classification in dermatology: A CNN-RF hybrid approach," in *Proc. 3rd Int. Conf. Artif. Intell. Internet Things (AllIoT)*, Vellore, India, May 2024, pp. 1–5.



- [17] S. Marulkar and B. Narain, "Nail disease prediction using a deep learning integrated framework," in *Proc. 3rd Int. Conf. Intell. Technol. (CONIT)*, Hubli, India, Jun. 2023, pp. 1–6.

การปรับปรุงประสิทธิภาพของต้นไม้เทอร์นารี เพื่อรองรับระดับการให้บริการที่มีคุณภาพแตกต่างกัน

วรกร ศรีเชวงทรัพย์¹ จิตติชญา ธนมิตรสมบุญ² กันติชา กิตติพิรชล^{3*}

^{1,2,3*} ห้องปฏิบัติการระบบการคำนวณขั้นสูง คณะวิศวกรรมศาสตร์ สถาบันเทคโนโลยีไทย-ญี่ปุ่น, กรุงเทพมหานคร, ประเทศไทย

*ผู้ประพันธ์บรรณกิจ อีเมล : kanticha@tni.ac.th

รับต้นฉบับ : 14 ตุลาคม 2568; รับบทความฉบับแก้ไข : 3 พฤศจิกายน 2568; ตอรับบทความ : 27 พฤศจิกายน 2568

เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

บทความนี้นำเสนอ 3 อัลกอริทึม ซึ่งปรับปรุงประสิทธิภาพของต้นไม้เทอร์นารีให้สามารถรองรับระดับคุณภาพการให้บริการที่แตกต่างกัน โดยอัลกอริทึมทั้ง 3 ได้แก่ อัลกอริทึม Partial Access ประเภทที่ 1 อัลกอริทึม Partial Access ประเภทที่ 2 และ อัลกอริทึม Adaptive Probability สำหรับอัลกอริทึมที่นำเสนอจะแบ่งผู้ใช้ออกเป็น 2 คลาส ได้แก่ คลาส 1 และ คลาส 2 กำหนดให้ผู้ใช้คลาส 1 มีลำดับความสำคัญสูงกว่าผู้ใช้คลาส 2 ในอัลกอริทึม Partial Access ประเภทที่ 1 ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 สามารถสุ่มเลือกได้ทั้ง 3 ช่องสัญญาณ อัลกอริทึมที่ 2 ได้แก่ อัลกอริทึม Partial Access ประเภทที่ 2 ในอัลกอริทึมนี้ ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณหลัง อัลกอริทึมที่ 3 ได้แก่ อัลกอริทึม Adaptive Probability ในอัลกอริทึมต้นไม้เทอร์นารี ผู้ใช้แต่ละคนจะสุ่มเลือก 1 ช่องสัญญาณ จากทั้งหมด 3 ช่องสัญญาณ ซึ่งเมื่อมองในรูปแบบของความน่าจะเป็น ผู้ใช้แต่ละรายสุ่มเลือกช่องสัญญาณแต่ละช่องด้วยความน่าจะเป็นเท่ากับ $1/3$ สำหรับอัลกอริทึม Adaptive Probability จะกำหนดให้ความน่าจะเป็นในการเข้าใช้ช่องสัญญาณแต่ละช่องไม่เท่ากัน อย่างเช่น กำหนดให้ความน่าจะเป็นของการเข้าใช้ช่องสัญญาณที่ 1 2 และ 3 เท่ากับ $1/2$ $2/5$ และ $1/10$ ตามลำดับ จากพฤติกรรมการเข้าถึงช่องสัญญาณที่ต่างกันระหว่างผู้ใช้คลาส 1 และผู้ใช้คลาส 2 ทำให้ผู้ใช้แต่ละคลาสมีค่าประวิงเวลาที่ไม่เท่ากัน จึงสามารถนำอัลกอริทึมทั้ง 3 นี้ มาใช้เพื่อรองรับระบบที่ต้องการคุณภาพการให้บริการที่ต่างกันได้ จากผลการทดสอบพบว่า ทุกอัลกอริทึมที่นำเสนอสามารถนำไปใช้ในระบบที่ต้องการรองรับคุณภาพการให้บริการที่ต่างกันได้ โดยเฉพาะอัลกอริทึม Adaptive Probability สามารถปรับค่าพารามิเตอร์ให้สามารถรองรับคุณภาพการให้บริการที่ต่างกันได้จำนวนหลายระดับ ในขณะที่ยังคงค่าประวิงเวลาที่เหมาะสม

คำสำคัญ: การปรับความน่าจะเป็น การเข้าใช้ช่องสัญญาณ การแก้ปัญหาการชนกัน การเข้าถึงบางส่วน อัลกอริทึมต้นไม้เทอร์นารี

Improving Efficiency of Ternary Tree to Support Different Quality of Service Levels

Warakorn Srichavengsup¹ Titichaya Thanamitsomboon² Kanticha Kittipeerachon^{3*}

^{1,2,3*} *Advanced Computation Systems Laboratory: ACS, Faculty of Engineering, Thai-Nichi Institute of Technology,
Bangkok, Thailand*

*Corresponding Author. E-mail address: kanticha@tni.ac.th

Received: 14 October 2025; Revised: 3 November 2025; Accepted: 27 November 2025

Published online: 26 December 2025

Abstract

This paper presents three algorithms that improve the performance of Ternary tree to support different quality of service levels. These algorithms are Partial Access type 1, Partial Access type 2 and Adaptive Probability algorithms. For the proposed algorithms, users are divided into two classes, namely class 1 and class 2, with class 1 users given higher priority than class 2 users. In Partial Access type 1 algorithm, class 1 users randomly select one slot from the first two slots, while class 2 users can access all three slots. In Partial Access type 2 algorithm, class 1 users randomly select 1 slot from the first 2 slots, while class 2 users randomly select 1 slot from the last 2 slots. Third algorithm is Adaptive Probability algorithm. In Ternary tree algorithm, each user randomly selects 1 slot out of 3 slots. When viewed in terms of probability, each user randomly accesses each slot with a probability of $1/3$. Adaptive probability algorithms use different probability values for each slot. For example, let the probability of accessing slots 1, 2, and 3 be $1/2$, $2/5$, and $1/10$, respectively. Due to the different channel access behavior between class 1 and class 2 users, each class of users has different delay values. Therefore, these three algorithms can be used to support systems that require different quality of service levels. The results show that each algorithm can provide different quality of service, especially Adaptive Probability algorithm, which can adjust its parameters to accommodate different quality of service levels while maintaining an appropriate delay.

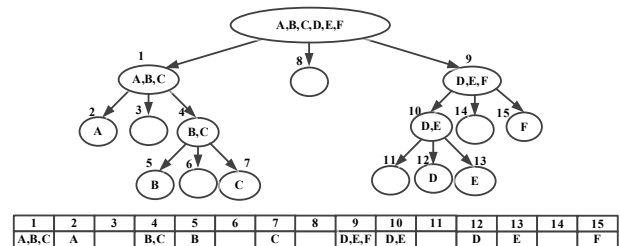
Keywords: Adaptive probability, Channel access, Collision resolution, Partial access, Ternary tree algorithm

1) บทนำ

ในปัจจุบัน ผู้ใช้งานเครือข่ายจำนวนมากมีความต้องการในการส่งข้อมูลจำนวนมาก เมื่อผู้ใช้หลายคนส่งข้อมูลพร้อมกัน จะทำให้เกิดการชนกันของข้อมูลขึ้น จึงมีความจำเป็นต้องมีอัลกอริทึมที่ใช้ควบคุมการเข้าถึงช่องสัญญาณเพื่อลดการชนกันของข้อมูล อัลกอริทึมที่นิยมใช้ในการควบคุมการเข้าถึงช่องสัญญาณแบ่งออกเป็น 2 กลุ่มหลัก ได้แก่ อัลกอริทึมที่ไม่มีการชนกันของข้อมูล (collision free algorithm) ตัวอย่างของอัลกอริทึมกลุ่มนี้ได้แก่ TDMA [1]–[3] FDMA [4]–[6] และ CDMA [7]–[9] โดยจะมีการแบ่งช่องสัญญาณให้ผู้ใช้งานแต่ละรายอย่างชัดเจน แต่อัลกอริทึมนี้ มีข้อจำกัดตรงที่ หากมีช่วงเวลาใดที่ผู้ใช้งานไม่มีข้อมูลที่ส่งช่องสัญญาณที่ว่างนั้น จะสามารถแบ่งให้ผู้ใช้งานอื่นใช้งานได้ และอัลกอริทึมนี้ไม่เหมาะที่จะรองรับผู้ใช้งานจำนวนมาก เพื่อให้รองรับจำนวนผู้ใช้งานได้มากขึ้น มีผู้เสนออัลกอริทึมที่มีการชนกันของข้อมูล (collision-based algorithm) ตัวอย่างของอัลกอริทึมกลุ่มนี้ได้แก่ อัลกอริทึมต้นไม้ไบนารี [10], [11] อัลกอริทึมต้นไม้เทอร์นารี [12] โดยอัลกอริทึมไบนารีจะแบ่งผู้ใช้ที่เกี่ยวข้องกับการชนกันออกเป็น 2 กลุ่ม ในขณะที่อัลกอริทึมเทอร์นารีจะแบ่งผู้ใช้ที่เกี่ยวข้องกับการชนกันออกเป็น 3 กลุ่ม และเมื่อเกิดการชนกันซ้ำจะมีการแบ่งกลุ่มย่อยซ้ำไปเรื่อย ๆ จนกว่าผู้ใช้ทุกคนจะประสบความสำเร็จในการเข้าใช้ช่องสัญญาณ เมื่อเปรียบเทียบสมรรถนะของอัลกอริทึมต้นไม้ไบนารีและอัลกอริทึมต้นไม้เทอร์นารี อัลกอริทึมต้นไม้เทอร์นารีจะมีสมรรถนะที่ดีกว่าโดยผู้ใช้จะมีค่าประวิงเวลาในการเข้าใช้ช่องสัญญาณน้อยกว่าอัลกอริทึมต้นไม้ไบนารี [12], [13]

บทความวิจัยฉบับนี้ มีวัตถุประสงค์เพื่อปรับปรุงอัลกอริทึมต้นไม้เทอร์นารีให้สามารถรองรับระดับการให้บริการที่มีคุณภาพแตกต่างกัน โดยนำเสนออัลกอริทึม 3 ตัว ได้แก่ อัลกอริทึม Partial Access ประเภทที่ 1 อัลกอริทึม Partial Access ประเภทที่ 2 และอัลกอริทึม Adaptive Probability โดยปกติในต้นไม้เทอร์นารี ผู้ใช้แต่ละคนจะสุ่มเลือก 1 ช่องสัญญาณ จากทั้งหมด 3 ช่องสัญญาณ ดังแสดงในรูปที่ 1 ผู้ใช้ทุกคนจะมีค่าประวิงเวลาโดยเฉลี่ยเท่ากัน สำหรับอัลกอริทึม Partial Access ประเภทที่ 1 ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 สามารถสุ่มเลือกได้ทั้ง 3 ช่องสัญญาณ อัลกอริทึมที่ 2 ได้แก่ อัลกอริทึม Partial Access ประเภทที่ 2 ในอัลกอริทึมนี้ ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 จะสุ่มเลือก 1 ช่องสัญญาณ

จาก 2 ช่องสัญญาณหลัง อัลกอริทึมที่ 3 ได้แก่ อัลกอริทึม Adaptive Probability ในอัลกอริทึมต้นไม้เทอร์นารี ผู้ใช้แต่ละคนจะสุ่มเลือก 1 ช่องสัญญาณ จากทั้งหมด 3 ช่องสัญญาณ ซึ่งเมื่อมองในรูปแบบของความน่าจะเป็น ผู้ใช้แต่ละรายสุ่มเข้าใช้ช่องสัญญาณแต่ละช่องด้วยความน่าจะเป็นเท่ากับ $1/3$ อัลกอริทึม Adaptive Probability นี้กำหนดให้ความน่าจะเป็นในการเข้าใช้ช่องสัญญาณแต่ละช่องไม่เท่ากัน โดยอัลกอริทึมที่นำเสนอสามารถนำไปใช้ในระบบเครือข่ายสื่อสารไร้สายในปัจจุบันได้อย่างเช่น เครือข่ายเซนเซอร์ไร้สาย เครือข่ายระบบอินเทอร์เน็ตในทุกสรรพสิ่ง และเครือข่ายยุค 5G เนื่องจากเครือข่ายข้างต้นจำเป็นต้องส่งข้อมูลหลากหลายประเภท ซึ่งข้อมูลแต่ละประเภทต้องการระดับคุณภาพการให้บริการที่แตกต่างกัน สำหรับรายละเอียดของทั้ง 3 อัลกอริทึมจะแสดงในหัวข้อถัดไป



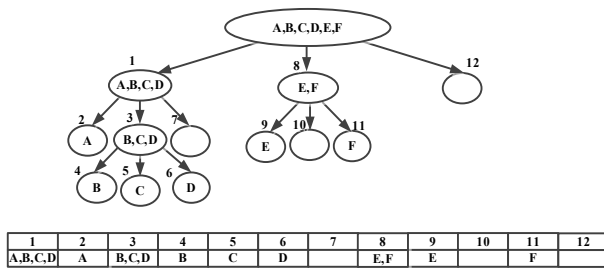
รูปที่ 1 : การแก้ปัญหาการชนกันของต้นไม้เทอร์นารี

2) อัลกอริทึมที่นำเสนอ

ในหัวข้อนี้ จะแสดงรายละเอียดของอัลกอริทึมที่นำเสนอ 3 ตัว โดยแสดงตัวอย่างประกอบ กำหนดให้ผู้ใช้ในระบบมีอยู่ 2 คลาส ได้แก่ ผู้ใช้คลาส 1 และผู้ใช้คลาส 2 โดยผู้ใช้คลาส 1 ต้องการระดับคุณภาพการให้บริการที่สูงกว่าผู้ใช้คลาส 2

2.1) อัลกอริทึม Partial Access ประเภทที่ 1

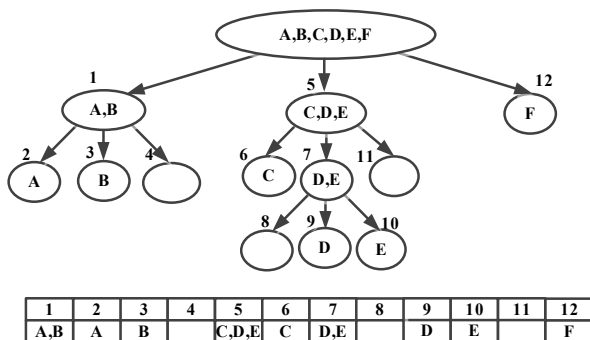
ในอัลกอริทึม Partial Access ประเภทที่ 1 ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 สามารถสุ่มเลือกได้ทั้ง 3 ช่องสัญญาณ ดังแสดงในรูปที่ 2 กำหนดให้ผู้ใช้ A B และ C เป็นผู้ใช้คลาส 1 และผู้ใช้ D E และ F เป็นผู้ใช้คลาส 2 จากรูปเมื่อมีการแบ่งกลุ่มเกิดขึ้นผู้ใช้ A B และ C จะสุ่มเลือก 1 ใน 2 ช่องสัญญาณแรกเสมอ ในขณะที่ผู้ใช้ D E และ F จะสุ่มเลือกช่องสัญญาณใดช่องสัญญาณหนึ่ง การดำเนินการเช่นนี้ จะทำให้ผู้ใช้คลาส 1 มีโอกาสประสบความสำเร็จในการเข้าใช้ช่องสัญญาณก่อนผู้ใช้คลาส 2 ส่งผลให้ผู้ใช้คลาส 1 มีค่าประวิงเวลาโดยเฉลี่ยต่ำกว่าผู้ใช้คลาส 2



รูปที่ 2 : การแก้ปัญหาการชนกันของอัลกอริทึม Partial Access ประเภทที่ 1

2.2) อัลกอริทึม Partial Access ประเภทที่ 2

อัลกอริทึมที่ 2 ได้แก่ อัลกอริทึม Partial Access ประเภทที่ 2 ในอัลกอริทึมนี้ ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 จะสุ่มเลือก 1 ช่องสัญญาณจาก 2 ช่องสัญญาณหลัง ดังแสดงในรูปที่ 3 จากรูปเมื่อมีการแบ่งกลุ่มเกิดขึ้น ผู้ใช้ A B และ C ซึ่งเป็นผู้ใช้คลาส 1 จะสุ่มเลือก 1 ใน 2 ช่องสัญญาณแรกเสมอ ในขณะที่ผู้ใช้ D E และ F ซึ่งเป็นผู้ใช้คลาส 2 จะสุ่มเลือก 1 ใน 2 ช่องสัญญาณหลังเสมอ ซึ่งจะทำให้ผู้ใช้คลาส 1 มีโอกาสประสบความสำเร็จในการเข้าใช้ช่องสัญญาณก่อนผู้ใช้คลาส 2 และผู้ใช้คลาส 1 มีค่าประสิทธิภาพโดยเฉลี่ยต่ำกว่าผู้ใช้คลาส 2



รูปที่ 3 : การแก้ปัญหาการชนกันของอัลกอริทึม Partial Access ประเภทที่ 2

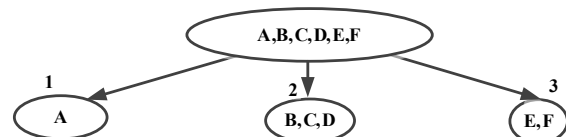
2.3) อัลกอริทึม Adaptive Probability

ในอัลกอริทึมต้นไม้เทอร์นารี ผู้ใช้แต่ละคนจะสุ่มเลือก 1 ช่องสัญญาณ จากทั้งหมด 3 ช่องสัญญาณ ซึ่งเมื่อมองในรูปแบบของความน่าจะเป็น ผู้ใช้แต่ละรายสุ่มเข้าใช้ช่องสัญญาณแต่ละช่องด้วยความน่าจะเป็นเท่ากับ $1/3$ อัลกอริทึม Adaptive Probability กำหนดให้ความน่าจะเป็นในการเข้าใช้ช่องสัญญาณแต่ละช่องไม่เท่ากัน อย่างเช่น กำหนดให้ความน่าจะเป็นในการเข้าใช้ช่องสัญญาณของผู้ใช้คลาส 1 เป็นดังนี้ ความน่าจะเป็นในการเข้า

ใช้ช่องสัญญาณที่ 1 เท่ากับ $1/2$ ความน่าจะเป็นของการเข้าใช้ช่องสัญญาณที่ 2 เท่ากับ $2/5$ และความน่าจะเป็นของการเข้าใช้ช่องสัญญาณที่ 3 เท่ากับ $1/10$ ดังแสดงในรูปที่ 4 โดยค่าความน่าจะเป็นทั้ง 3 ข้างต้น เป็นค่าความน่าจะเป็นที่สมมติขึ้นมา เพื่อให้การยกตัวอย่างเห็นภาพที่ชัดเจนขึ้น ซึ่งจะพบว่าโอกาสในการเข้าใช้ช่องสัญญาณ 1 เท่ากับ 0.5 โอกาสในการเข้าใช้ช่องสัญญาณ 1 และ 2 เท่ากับ 0.9 ในขณะที่อัลกอริทึมต้นไม้เทอร์นารีมีโอกาสในการเข้าใช้ช่องสัญญาณ 1 เท่ากับ 0.33 โอกาสในการเข้าใช้ช่องสัญญาณ 1 และ 2 เท่ากับ 0.66 จากรูป ผู้ใช้ A ซึ่งเป็นผู้ใช้คลาส 1 สุ่มค่าความน่าจะเป็นเท่ากับ 0.38 ซึ่งมีค่าน้อยกว่า 0.5 ผู้ใช้ A จึงสุ่มเลือกช่องสัญญาณแรก ผู้ใช้ B ซึ่งเป็นผู้ใช้คลาส 1 สุ่มค่าความน่าจะเป็นเท่ากับ 0.68 ซึ่งมีค่ามากกว่า 0.5 แต่น้อยกว่า 0.9 ผู้ใช้ B จึงสุ่มเลือกช่องสัญญาณที่ 2 ผู้ใช้รายอื่น ๆ จะใช้หลักการเดียวกันในการสุ่มเลือกช่องสัญญาณ

การกำหนดค่าความน่าจะเป็นของผู้ใช้คลาส 1 ในการเข้าถึงช่องสัญญาณที่ 1 (P_1) และค่าความน่าจะเป็นของผู้ใช้คลาส 1 ในการเข้าถึงช่องสัญญาณที่ 1 และ 2 (P_2) ตามตัวอย่าง จะทำให้ผู้ใช้คลาส 1 มีโอกาสประสบความสำเร็จในการเข้าใช้ช่องสัญญาณก่อนผู้ใช้คลาส 2

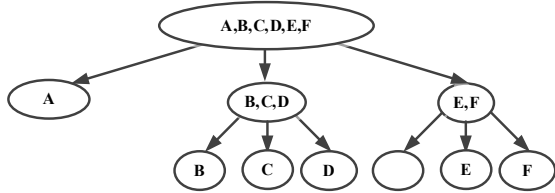
Class	1	1	1	2	2	2
User	A	B	C	D	E	F
P ₁	0.5	0.5	0.5	0.33	0.33	0.33
P ₂	0.9	0.9	0.9	0.66	0.66	0.66
rand	0.38	0.68	0.8	0.36	0.69	0.70
slot	1	2	2	2	3	3



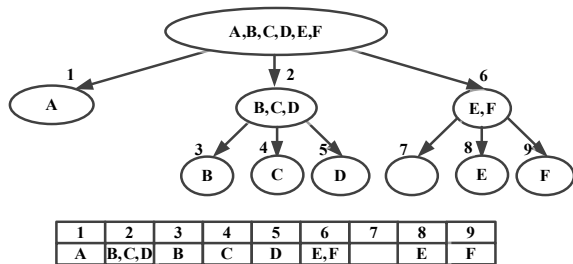
รูปที่ 4 : การแก้ปัญหาการชนกันของอัลกอริทึม Adaptive Probability ขั้นที่ 1

จากรูปที่ 4 จะพบว่า ผู้ใช้ B C และ D สุ่มเลือกช่องสัญญาณที่ 2 เหมือนกัน ทำให้เกิดการชนกันขึ้น ในขณะที่ผู้ใช้ E และ F ก็สุ่มเลือกช่องสัญญาณเดียวกัน จึงต้องมีการแบ่งกลุ่มครั้งที่ 2 ดังแสดงในรูปที่ 5 หากยังเกิดการชนกันขึ้นอีกจะมีการแบ่งกลุ่มต่อไปเรื่อย ๆ จนกว่าผู้ใช้ทุกคนประสบความสำเร็จในการเข้าใช้ช่องสัญญาณ รูปที่ 6 แสดงหมายเลขช่องสัญญาณที่ผู้ใช้แต่ละรายเข้าใช้

Class	1	1	2	2	2
User	B	C	D	E	F
P ₁	0.5	0.5	0.33	0.33	0.33
P ₂	0.9	0.9	0.66	0.66	0.66
rand	0.40	0.61	0.7	0.36	0.69
slot	1	2	3	2	3



รูปที่ 5 : การแก้ปัญหาการชนกันของอัลกอริทึม Adaptive Probability ขั้นที่ 2



รูปที่ 6 : การแก้ปัญหาการชนกันของอัลกอริทึม Adaptive Probability

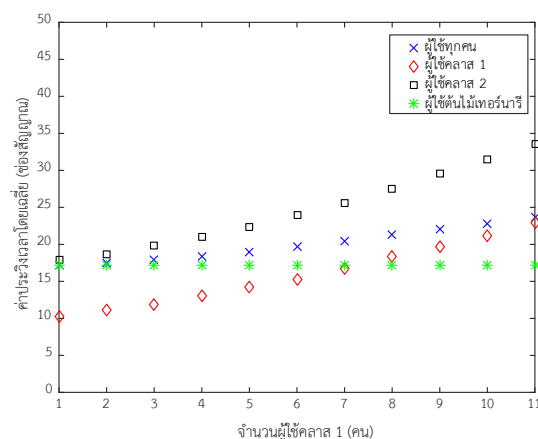
จากรายละเอียดของอัลกอริทึมทั้ง 3 ข้างต้น โอกาสในการประสบความสำเร็จในการเข้าใช้ช่องสัญญาณของผู้ใช้คลาส 1 ต่างจากโอกาสในการประสบความสำเร็จในการเข้าใช้ช่องสัญญาณของผู้ใช้คลาส 2 จึงสามารถนำอัลกอริทึมที่เสนอ มาใช้รองรับ ทราฟฟิกที่มีความต้องการคุณภาพการให้บริการที่แตกต่างกันได้

3) ผลการวิจัยและการอภิปราย

ในหัวข้อนี้ จะแสดงสมรรถนะและเปรียบเทียบอัลกอริทึมที่นำเสนอ กำหนดให้ค่าดัชนีคุณภาพการให้บริการ (QoS index) เท่ากับ อัตราส่วนระหว่างค่าประวิงเวลาโดยเฉลี่ยระหว่างผู้ใช้คลาส 2 และผู้ใช้คลาส 1 หากค่าดัชนีคุณภาพการให้บริการมีค่ามากแสดงว่า มีความแตกต่างระหว่างค่าประวิงเวลาโดยเฉลี่ยระหว่างผู้ใช้คลาส 2 และผู้ใช้คลาส 1 มาก

รูปที่ 7 แสดงค่าประวิงเวลาโดยเฉลี่ยเมื่อปรับจำนวนผู้ใช้คลาส 1 ของอัลกอริทึม Partial Access ประเภทที่ 1 กำหนดให้จำนวนผู้ใช้ทั้งหมดเท่ากับ 12 คน สำหรับอัลกอริทึมนี้ ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณ จาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 สุ่มเลือก 1 ช่องสัญญาณ จากช่องสัญญาณทั้งหมด 3 ช่อง จากรูปจะพบว่า ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2

สูงขึ้นตามจำนวนผู้ใช้คลาส 1 เนื่องจากเมื่อผู้ใช้คลาส 1 มีจำนวนมากขึ้นโอกาสที่จะเกิดการชนกันระหว่างผู้ใช้คลาส 1 ด้วยกันเอง และระหว่างผู้ใช้คลาส 1 และ 2 ใน 2 ช่องสัญญาณแรกจะมากขึ้นตาม และทำให้ผู้ใช้คลาส 1 และ 2 มีโอกาสประสบความสำเร็จในการจองช่องสัญญาณแรก ๆ ลดลง นอกจากนี้จะพบว่า เมื่อผู้ใช้คลาส 1 เท่ากับ 1 คน ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ทุกคนจะใกล้เคียงกับค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 เนื่องจากผู้ใช้ส่วนใหญ่เป็นผู้ใช้คลาส 2 และเมื่อผู้ใช้คลาส 1 เท่ากับ 11 คน จะพบว่าค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ทุกคนจะใกล้เคียงกับค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 เนื่องจากผู้ใช้ส่วนใหญ่เป็นผู้ใช้คลาส 1 ทั้งนี้เมื่อเปรียบเทียบกับอัลกอริทึมต้นไม้เทอร์นารีพบว่าเมื่อจำนวนผู้ใช้คลาส 1 เท่ากับ 1 คน ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 และผู้ใช้ทุกคนของอัลกอริทึม Partial Access ประเภทที่ 1 จะใกล้เคียงกับค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ของอัลกอริทึมต้นไม้เทอร์นารี เนื่องจากการที่มีผู้ใช้คลาส 1 เพียงคนเดียวสุ่มเข้าใช้ 2 ช่องสัญญาณแรก การดำเนินการเช่นนี้ส่งผลกระทบต่อค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 น้อย แต่เมื่อผู้ใช้คลาส 1 มีจำนวนมากขึ้น ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2 จะสูงขึ้น เนื่องจากการที่มีผู้ใช้คลาส 1 จำนวนมากจะทำให้มีโอกาสสูงที่จะเกิดการชนกันเองระหว่างผู้ใช้คลาส 1 และการชนกันระหว่างผู้ใช้คลาส 1 และ 2 ใน 2 ช่องสัญญาณแรก ซึ่งส่งผลให้ค่าประวิงเวลาโดยเฉลี่ยสูงขึ้น โดยมีค่ามากกว่าค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ของอัลกอริทึมต้นไม้เทอร์นารี นั้นแสดงให้เห็นว่า อัลกอริทึมที่นำเสนอสามารถปรับปรุงอัลกอริทึมต้นไม้เทอร์นารี ให้สามารถรองรับระดับการให้บริการที่มีคุณภาพที่แตกต่างกันได้ แต่ต้องแลกกับค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ที่เพิ่มขึ้น

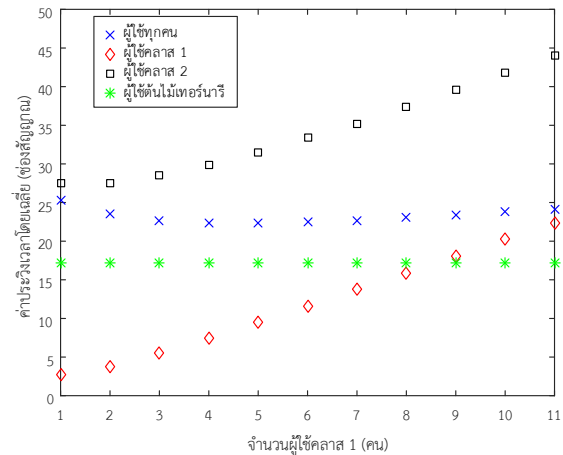


รูปที่ 7 : ค่าประวิงเวลาโดยเฉลี่ยเมื่อปรับจำนวนผู้ใช้คลาส 1 ของอัลกอริทึม Partial Access ประเภทที่ 1

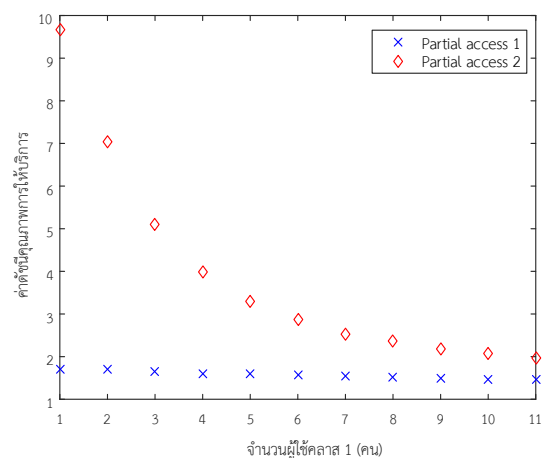
รูปที่ 8 แสดงค่าประวิงเวลาโดยเฉลี่ยเมื่อปรับจำนวนผู้ใช้คลาส 1 ของอัลกอริทึม Partial Access ประเภทที่ 2 กำหนดให้จำนวนผู้ใช้ทั้งหมดเท่ากับ 12 คน สำหรับอัลกอริทึมนี้ ผู้ใช้คลาส 1 จะสุ่มเลือก 1 ช่องสัญญาณ จาก 2 ช่องสัญญาณแรก ในขณะที่ผู้ใช้คลาส 2 สุ่มเลือก 1 ช่องสัญญาณ จาก 2 ช่องสัญญาณหลังจากที่ผู้ใช้คลาส 1 สุ่มเลือก 1 ช่องสัญญาณแล้ว จากรูปจะพบว่า ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2 สูงขึ้นตามจำนวนผู้ใช้คลาส 1 โดยมีเหตุการณ์การเพิ่มของค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2 เช่นเดียวกับผลของอัลกอริทึม Partial Access ประเภทที่ 1 แต่สิ่งที่แตกต่างกันคือในอัลกอริทึม Partial Access ประเภทที่ 2 ค่าประวิงเวลาโดยเฉลี่ยระหว่างผู้ใช้คลาส 1 และผู้ใช้คลาส 2 ต่างกันมาก เนื่องจากในอัลกอริทึม Partial Access ประเภทที่ 1 ผู้ใช้คลาส 2 เข้าใช้ได้ทุกช่องสัญญาณ แต่ในอัลกอริทึม Partial Access ประเภทที่ 2 ผู้ใช้คลาส 2 เข้าถึงได้เฉพาะ 2 ช่องสัญญาณหลัง ทำให้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 สูงขึ้น และเมื่อช่องสัญญาณแรกมีเฉพาะผู้ใช้คลาส 1 เท่านั้นที่เข้าถึงได้ ผู้ใช้คลาส 1 จึงมีโอกาสประสบความสำเร็จในช่องสัญญาณแรก ๆ มากขึ้น ทำให้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 มีค่าลดลง ทั้งนี้เมื่อเปรียบผลกับอัลกอริทึมต้นไม้เทอร์นารี พบว่าเมื่อผู้ใช้คลาส 1 มีจำนวนมากขึ้น ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2 จะสูงขึ้น เนื่องจากการที่มีผู้ใช้คลาส 1 จำนวนมากจะทำให้มีโอกาสสูงที่จะเกิดการชนกันเองระหว่างผู้ใช้คลาส 1 ใน 2 ช่องสัญญาณแรก และมีการชนกันระหว่างผู้ใช้คลาส 1 ใน 2 ในช่องสัญญาณที่ 2 ซึ่งส่งผลให้ค่าประวิงเวลาโดยเฉลี่ยสูงขึ้น โดยมีค่ามากกว่าค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ของอัลกอริทึมต้นไม้เทอร์นารี อัลกอริทึมที่นำเสนอสามารถปรับปรุงอัลกอริทึมต้นไม้เทอร์นารี ให้สามารถรองรับระดับการให้บริการที่มีคุณภาพที่แตกต่างกันได้ แต่ต้องแลกกับค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้ที่เพิ่มขึ้น

รูปที่ 9 แสดงการเปรียบเทียบค่าดัชนีคุณภาพการให้บริการเมื่อปรับจำนวนผู้ใช้คลาส 1 จากรูปจะพบว่า ค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Partial Access ประเภทที่ 1 มีค่ามากกว่า 1 แต่น้อยกว่า 2 และค่านี้มีแนวโน้มค่อนข้างคงที่เนื่องจากอัตราการเพิ่มค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 และ 1 ค่อนข้างใกล้เคียงกัน ดังแสดงในรูปที่ 7 ในขณะที่ ค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Partial Access ประเภทที่ 2 มีค่าสูงที่ค่าจำนวนผู้ใช้คลาส 1 เท่ากับ 1 เนื่องจากในกรณีที่จำนวนผู้ใช้คลาส 2 เท่ากับ 11 คน และผู้ใช้ทั้ง 11 คนแย่งเข้าใช้ 2 ช่องสัญญาณหลัง ทำให้มีโอกาสชนกันระหว่างผู้ใช้คลาส 2 สูง

ส่งผลให้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 สูงตามไปด้วย เมื่อจำนวนผู้ใช้คลาส 1 มีจำนวนมากขึ้น ค่าดัชนีคุณภาพการให้บริการจะมีค่าลดลง เนื่องจาก อัตราการเพิ่มของค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 มีค่าน้อยกว่าอัตราการเพิ่มของค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 ดังแสดงในรูปที่ 8



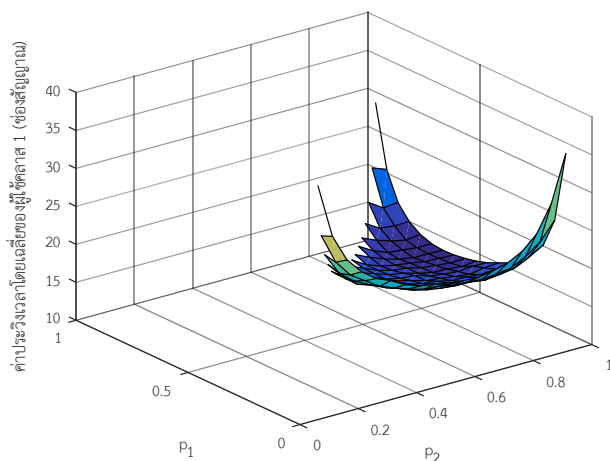
รูปที่ 8 : ค่าประวิงเวลาโดยเฉลี่ยเมื่อปรับจำนวนผู้ใช้คลาส 1 ของอัลกอริทึม Partial Access ประเภทที่ 2



รูปที่ 9 : การเปรียบเทียบค่าดัชนีการให้บริการเมื่อปรับจำนวนผู้ใช้คลาส 1 ของอัลกอริทึม Partial Access ประเภทที่ 1 และ 2

รูปที่ 10 แสดงค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 ของอัลกอริทึม Adaptive Probability เมื่อปรับค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 (P_1) และค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 (P_2) กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 ที่เหมาะสมได้แก่ ค่าประวิงเวลาโดยเฉลี่ยซึ่งมีค่าต่ำที่สุดเพราะจะเป็นค่าที่ผู้ใช้คลาส 1 ใช้เวลารอในการเข้าใช้ช่องสัญญาณน้อย

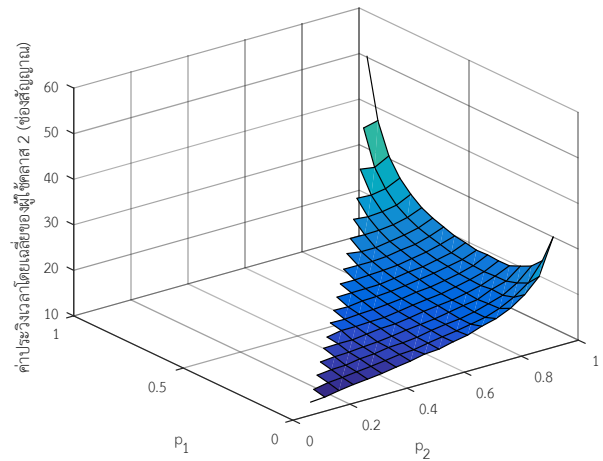
ที่สุด โดยมีค่าประวิงเวลาโดยเฉลี่ยเท่ากับ 14.72 ช่องสัญญาณ ซึ่งเกิดขึ้นที่ค่า P_1 และ P_2 เท่ากับ 0.60 และ 0.95 ตามลำดับ เนื่องจากที่ค่าความน่าจะเป็นชุดนี้ ผู้ใช้คลาส 1 มีโอกาสสูงที่จะเข้าใช้ช่องสัญญาณที่ 1 และ 2 และมีโอกาสชนกันเองระหว่างผู้ใช้คลาส 1 น้อย ในกรณีที่ค่า P_1 และ P_2 มีค่าน้อยเกินไป จะทำให้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 สูง ยกตัวอย่างเช่น กรณีที่ค่า P_1 และ P_2 เท่ากับ 0.05 และ 0.1 ตามลำดับ กรณีนี้ผู้ใช้คลาส 1 มีโอกาสเข้าใช้ช่องสัญญาณที่ 3 สูง ทำให้จำนวนครั้งในการชนกันระหว่างผู้ใช้คลาส 1 ในช่องสัญญาณที่ 3 สูง และกรณีที่ค่า P_1 และ P_2 เท่ากับ 0.05 และ 0.95 ตามลำดับ กรณีนี้ผู้ใช้คลาส 1 มีโอกาสเข้าใช้ช่องสัญญาณที่ 2 สูง ทำให้จำนวนครั้งในการชนกันระหว่างผู้ใช้คลาส 1 ในช่องสัญญาณที่ 2 สูง ซึ่งค่า P_1 ควรจะมีค่าปานกลาง และค่า P_2 ควรจะมีค่าสูง เพื่อให้ผู้ใช้คลาส 1 จะได้มีโอกาสประสบความสำเร็จในการเข้าใช้ช่องสัญญาณแรก ๆ สูง



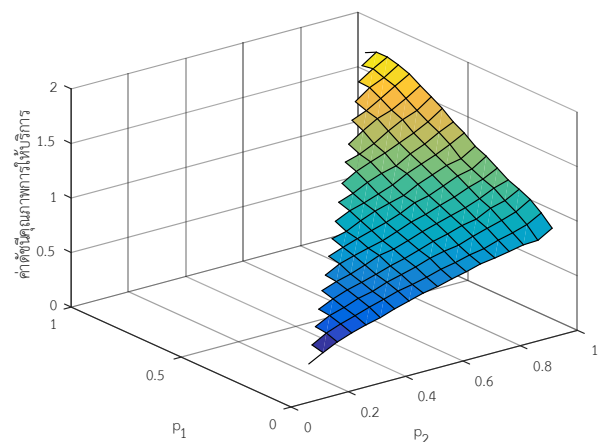
รูปที่ 10 : ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 ของอัลกอริทึม Adaptive Probability เมื่อปรับค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 (P_1) และ ค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 (P_2) กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน

รูปที่ 11 แสดงค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 ของอัลกอริทึม Adaptive Probability เมื่อปรับค่า P_1 และ P_2 กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 ที่เหมาะสมได้แก่ ค่าประวิงเวลาโดยเฉลี่ยซึ่งมีค่าต่ำที่สุดเพราะจะเป็นค่าที่ผู้ใช้คลาส 2 ใช้เวลารอในการเข้าใช้ช่องสัญญาณน้อยที่สุด โดยมีค่าประวิงเวลาโดยเฉลี่ยเท่ากับ 11.06 ช่องสัญญาณ ซึ่งเกิดขึ้นที่ค่า P_1 และ P_2 เท่ากับ 0.05 และ 0.10 ตามลำดับ เนื่องจากที่ค่าความน่าจะเป็นชุดนี้

ผู้ใช้คลาส 1 มีโอกาสเข้าใช้ช่องสัญญาณที่ 3 สูง ทำให้ผู้ใช้คลาส 2 มีโอกาสประสบความสำเร็จในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 สูง ส่งผลให้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 ต่ำ



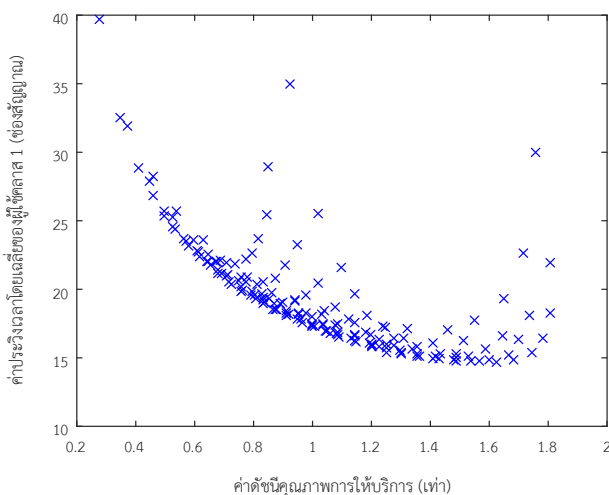
รูปที่ 11 : ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 ของอัลกอริทึม Adaptive Probability เมื่อปรับค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 (P_1) และ ค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 (P_2) กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน



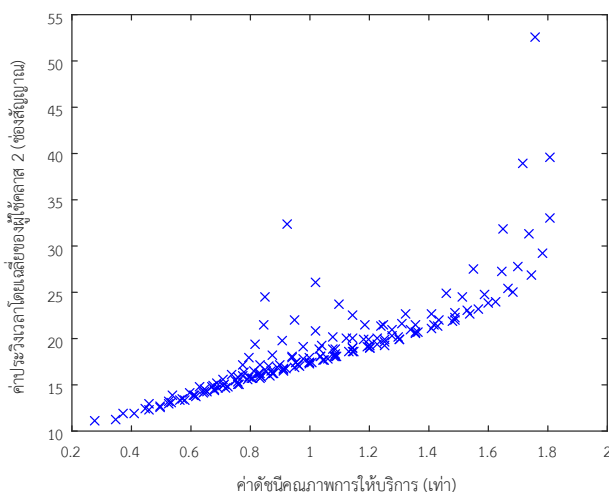
รูปที่ 12 : ค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Adaptive Probability เมื่อปรับค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 (P_1) และ ค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 (P_2) กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน

รูปที่ 12 แสดงค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Adaptive Probability เมื่อปรับค่า P_1 และ P_2 กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน จากรูปจะพบว่าค่าดัชนีคุณภาพการให้บริการมีค่าต่ำสุดที่ค่า P_1 และ P_2 เท่ากับ 0.05 และ 0.10 ตามลำดับ เนื่องจากที่ค่านี้นี้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส

1 มีค่าสูง และค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 มีค่าต่ำที่สุด และค่าดัชนีคุณภาพการให้บริการมีค่าสูงสุดที่ค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 มีค่ามากที่สุดที่ค่า P_1 และ P_2 เท่ากับ 0.9 และ 0.95 ตามลำดับ เนื่องจากกรณีนี้ผู้ใช้คลาส 1 มีโอกาสเข้าใช้ช่องสัญญาณที่ 1 สูง ทำให้มีโอกาสชนกันระหว่างผู้ใช้คลาส 1 และผู้ใช้คลาส 2 สูง ในช่องสัญญาณที่ 1 และ มีโอกาสสูงที่จะเกิดการชนกันซ้ำไปเรื่อย ๆ ทำให้ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 มีค่าสูงตามไปด้วย

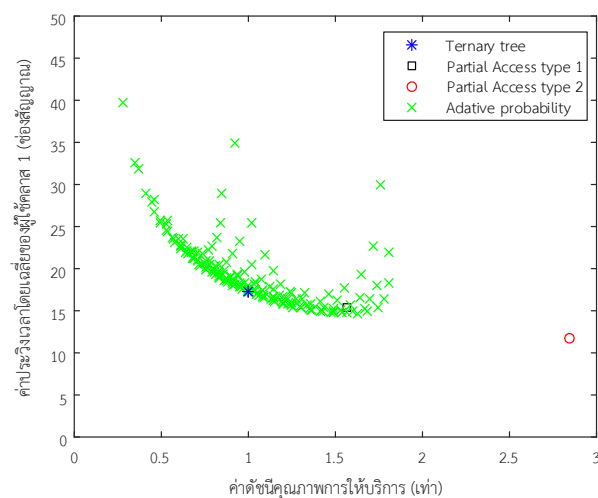


รูปที่ 13 : ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 เมื่อเทียบกับค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Adaptive Probability กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน

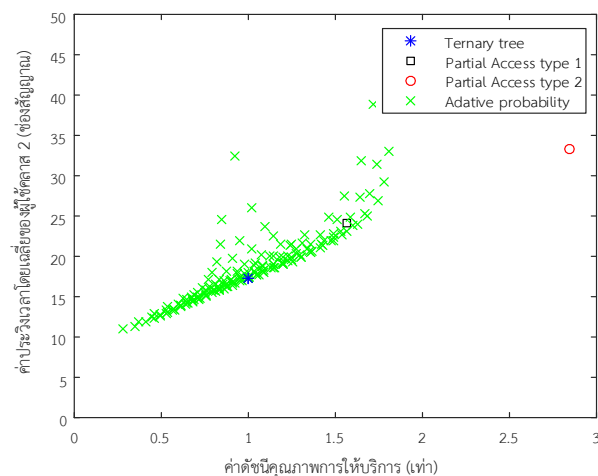


รูปที่ 14 : ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 เมื่อเทียบกับค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Adaptive Probability กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน

รูปที่ 13 และ 14 แสดงค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2 ที่ค่าดัชนีคุณภาพการให้บริการต่าง ๆ ของอัลกอริทึม Adaptive Probability จากรูปจะพบว่าที่ค่าดัชนีคุณภาพการให้บริการของอัลกอริทึม Adaptive Probability สูง ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 มีค่าต่ำ ในขณะที่ค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 มีค่าสูง นอกจากนี้ยังพบว่าอัลกอริทึม Adaptive Probability ให้ค่าดัชนีคุณภาพการให้บริการที่แตกต่างกันเป็นจำนวนมาก ในขณะที่ยังคงค่าประวิงเวลาที่เหมาะสม



รูปที่ 15 : การเปรียบเทียบค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 เมื่อเทียบกับค่าดัชนีคุณภาพการให้บริการของทุกอัลกอริทึม กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน



รูปที่ 16: การเปรียบเทียบค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 2 เมื่อเทียบกับค่าดัชนีคุณภาพการให้บริการของทุกอัลกอริทึม กำหนดจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน

รูปที่ 15 และ 16 แสดงการเปรียบเทียบค่าประวิงเวลาโดยเฉลี่ยของผู้ใช้คลาส 1 และ 2 ตามลำดับ เมื่อเทียบกับค่าดัชนีคุณภาพการให้บริการของทุกอัลกอริทึม กำหนดให้จำนวนผู้ใช้ทั้งหมดเท่ากับ 12 คน และจำนวนผู้ใช้คลาส 1 และ 2 เท่ากับ 6 คน โดยอัลกอริทึมที่นำมาเปรียบเทียบได้แก่ ต้นไม้เทอร์นารี อัลกอริทึม Partial Access ประเภทที่ 1 อัลกอริทึม Partial Access ประเภทที่ 2 และอัลกอริทึม Adaptive Probability จากรูปจะพบว่า อัลกอริทึมต้นไม้เทอร์นารีมีค่าดัชนีคุณภาพการให้บริการเท่ากับ 1 เนื่องจากผู้ใช้ทุกรายมีความเท่าเทียมกันในการใช้งานช่องสัญญาณ นอกจากนี้จะพบว่าอัลกอริทึม Adaptive Probability มีค่าดัชนีคุณภาพการให้บริการต่าง ๆ เป็นจำนวนมาก เนื่องจากอัลกอริทึมนี้สามารถปรับค่าพารามิเตอร์ 2 ตัว ได้แก่ ค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 (P_1) และค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 (P_2) ในขณะที่ต้นไม้เทอร์นารี อัลกอริทึม Partial Access ประเภทที่ 1 และอัลกอริทึม Partial Access ประเภทที่ 2 แต่ละอัลกอริทึมให้ค่าดัชนีคุณภาพการให้บริการเพียงค่าเดียว ดังนั้นจึงสรุปได้ว่า อัลกอริทึม Adaptive Probability มีความยืดหยุ่นสูง ให้ค่าประวิงเวลาที่เหมาะสม และสามารถนำมาใช้ในระบบที่ต้องการคุณภาพการให้บริการที่แตกต่างกันหลายระดับ

4) สรุปผล

บทความวิจัยนี้มีวัตถุประสงค์เพื่อปรับปรุงประสิทธิภาพของอัลกอริทึมต้นไม้เทอร์นารีให้สามารถรองรับระดับคุณภาพการให้บริการที่แตกต่างกัน โดยเสนออัลกอริทึม Partial Access ประเภทที่ 1 ซึ่งกำหนดให้ผู้ใช้คลาส 1 เข้าใช้ 2 ช่องสัญญาณแรกเท่านั้น อัลกอริทึมที่ 2 ได้แก่ อัลกอริทึม Partial Access ประเภทที่ 2 ซึ่งจำกัดการเข้าใช้ช่องสัญญาณของผู้ใช้คลาส 1 และ 2 และอัลกอริทึมที่ 3 ได้แก่ อัลกอริทึม Adaptive Probability ซึ่งใช้ค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 (P_1) และค่าความน่าจะเป็นในการเข้าใช้ช่องสัญญาณที่ 1 และ 2 (P_2) ในการควบคุมการเข้าใช้ช่องสัญญาณของผู้ใช้คลาส 1 โดยทั้ง 3 อัลกอริทึม ผู้ใช้คลาส 1 และผู้ใช้คลาส 2 มีพฤติกรรมการเข้าถึงช่องสัญญาณที่แตกต่างกัน ทำให้ผู้ใช้แต่ละคลาสมีค่าประวิงเวลาที่ไม่เท่ากัน จึงสามารถนำอัลกอริทึมทั้ง 3 นี้ มาใช้เพื่อรองรับระบบที่ต้องการคุณภาพการให้บริการที่แตกต่างกันได้ จากผลการทดสอบพบว่า อัลกอริทึม Adaptive Probability มี

ความยืดหยุ่นสูง โดยให้ค่าดัชนีคุณภาพการให้บริการที่แตกต่างกันเป็นจำนวนมาก ในขณะที่ยังคงค่าประวิงเวลาที่เหมาะสม จึงเหมาะที่จะนำมาให้บริการในระบบที่ต้องการคุณภาพการให้บริการที่แตกต่างกันหลายระดับ

REFERENCES

- [1] D. Niyato, P. Wang, and D. I. Kim, "Performance analysis and optimization of TDMA Network with wireless energy transfer," *IEEE Trans. Wirel. Commun.*, vol. 13, no. 8, pp. 4205–4219, 2014.
- [2] G. Pierobon, A. Zanella, and A. Salloum, "Contention-TDMA protocol: Performance evaluation," *IEEE Trans. Veh. Technol.*, vol. 51, no. 4, pp. 781–788, 2002.
- [3] X. Wang, A. G. Marques, and G. B. Giannakis, "Power-efficient resource allocation and quantization for TDMA using adaptive transmission and limited-rate feedback," *IEEE Trans. Signal Process.*, vol. 56, no. 9, pp. 4470–4485, 2008.
- [4] J. Zhang, L. L. Yang, L. Hanzo, and H. Gharavi, "Advances in cooperative single-carrier FDMA communications: Beyond LTE-advanced," *IEEE Commun. Surv. Tutor.*, vol. 17, no. 2, pp. 730–756, 2015.
- [5] F. R. Farrokhi, A. Lozano, G. J. Foschini and R. A. Valenzuela, "Spectral efficiency of FDMA/TDMA wireless systems with transmit and receive antenna arrays," *IEEE Trans. Wirel. Commun.*, vol. 1, no. 4, pp. 591–599, 2002.
- [6] H. G. Myung, J. Lim, and D. J. Goodman, "Single carrier FDMA for uplink wireless transmission," *IEEE Veh. Technol. Mag.*, vol. 1, no. 3, pp. 30–38, 2006.
- [7] S. Hara and R. Prasad, "Overview of multicarrier CDMA," *IEEE Commun. Mag.*, vol. 35, no. 12, pp. 126–133, 1997.
- [8] S. M. Alamouti, "A simple transmit diversity technique for wireless communications," *IEEE J. Sel. Areas Commun.*, vol. 16, no. 8, pp. 1451–1458, 1998.
- [9] X. Peng, K. -B. Png, Z. Lei, F. Chin, and C. C. Ko, "Two-layer spreading CDMA: An improved method for broadband uplink transmission," *IEEE Trans. Veh. Technol.*, vol. 57, no. 6, pp. 3563–3577, 2008.
- [10] Y. Y. Guo, H. W. Ding, Y. F. Zhao, J. Guo, and Q. L. Liu, "Probability detection CSMA with monitoring based on the improved binary tree conflict resolution algorithm," *Appl. Mech. Mater.*, vol. 610, pp. 897–904, 2014.



- [11] W. Srichavengsup and K. Kittipeerachon, "Performance evaluation of modified tree algorithm for supporting traffic with different priority requirements," *J. Eng. Digit. Technol. (JEDT)*, vol. 10, no. 2, pp. 46–56, 2022.
- [12] P. Mathys and P. Flajolet, "Q-ary collision resolution algorithms in random-access systems with free or blocked channel access," *IEEE Trans. Inf. Theory*, vol. 31, no. 2, pp. 217–243, 1985.
- [13] H. Wu and Y. Pan, *Medium Access Control in Wireless Networks*. New York, NY, USA: Nova Science, 2008.

การปรับปรุงประสิทธิภาพแกวคอยานพาหนะและการลำเลียงสินค้า ผ่านการจำลองสถานการณ์ : กรณีศึกษา

สุริพร ยืนยงค์¹ อธิธณณัฐ ปลุกผึ้ง² บุณย์สิริวรรธ ขำทัพ³ อนนจ ชัยมณี^{4*}

^{1,2,3,4*} ภาควิชาวิศวกรรมอุตสาหการ คณะวิศวกรรมศาสตร์ กำแพงแสน มหาวิทยาลัยเกษตรศาสตร์ วิทยาเขตกำแพงแสน,
นครปฐม, ประเทศไทย

*ผู้ประพันธ์บรรณกิจ อีเมล : anot.c@ku.th

รับต้นฉบับ : 14 พฤษภาคม 2568; รับบทความฉบับแก้ไข : 26 ตุลาคม 2568; ตอบรับบทความ : 20 พฤศจิกายน 2568
เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

งานวิจัยนี้ศึกษากระบวนการเพื่อเพิ่มประสิทธิภาพการจัดส่งสินค้า วัตถุประสงค์เพื่อลดเวลารอคอยของรถขนส่ง ลดเวลาการลำเลียงสินค้าภายในคลัง ส่งผลให้รอบการจัดส่งเพิ่มขึ้น การจำลองสถานการณ์ร่วมกับการจัดระเบียบแกวคอยานพาหนะและทฤษฎีการจัดการคลังสินค้าถูกนำมาใช้ในงานวิจัย เริ่มจากเก็บรวบรวมข้อมูลจากกระบวนการทำงานปัจจุบันมาสร้างเป็นแบบจำลองสถานการณ์ทางคอมพิวเตอร์โดยใช้โปรแกรม FlexSim เพื่อใช้วิเคราะห์ปัญหาและหาสถานการณ์เพื่อปรับปรุงกระบวนการ งานวิจัยเสนอวิธีการจัดการแกวคอย 3 สถานการณ์ การจัดการคลังสินค้าจากการวิเคราะห์ FSN และการจัดการแกวคอยร่วมกับการวิเคราะห์ FSN ทำให้ได้สถานการณ์ปรับปรุงทั้งหมด 7 สถานการณ์ ผลการจำลองสถานการณ์พบว่าการจัดแกวคอยที่ 1 การให้ความสำคัญกับรถหลักพร้อมกับการวิเคราะห์ FSN เป็นแนวทางการปรับปรุงกระบวนการที่เหมาะสมที่สุด วิธีการดังกล่าวทำให้ประสิทธิภาพของกระบวนการเพิ่มขึ้นเมื่อเปรียบเทียบกับสถานการณ์ก่อนปรับปรุง นั่นคือสามารถลดเวลารอคอยเฉลี่ยของรถทุกประเภทและทำให้จำนวนรถทั้งหมดที่ออกจากระบบโดยเฉลี่ยในช่วงเวลาที่กำหนดตลอดระยะเวลา 1 เดือน เพิ่มขึ้นมากที่สุดจาก 582 คัน เป็น 658 คัน จำนวนรถที่เพิ่มขึ้นนี้สามารถเพิ่มโอกาสในการขายสินค้า ส่งผลให้กำไรของบริษัทกรณีศึกษาเพิ่มขึ้น 14,063,037.11 บาท/เดือน

คำสำคัญ : การวิเคราะห์ FSN การจัดระเบียบแกวคอย การจำลองสถานการณ์



Improving the Efficiency of Vehicle Queuing and Product Conveyance through Simulation: A Case Study

Sureeporn Yeanyong¹ Aitnanat Plukfung² Boonsiwatt Khuamthap³ Anot Chaimanee^{4*}

^{1,2,3,4*}*Department of Industrial Engineering, Faculty of Engineering at Kamphaeng Saen,
Kasetsart University Kamphaeng Saen Campus, Nakhon Pathom, Thailand*

*Corresponding Author. E-mail address: anot.c@ku.th

Received: 14 May 2025; Revised: 26 October 2025; Accepted: 20 November 2025

Published online: 26 December 2025

Abstract

This research examines the process of enhancing product delivery efficiency. The objective is to reduce both vehicle waiting times and product conveying times in the warehouse, thereby increasing the overall number of deliveries. Simulation techniques were applied in combination with vehicle queuing disciplines, and warehouse management theory was incorporated into the research. The research, starting from data from the current process of the case study company, was collected and used to develop a computer simulation model in FlexSim, enabling both problem analysis and the exploration of improvement scenarios. Three queuing discipline scenarios, warehouse management from FSN analysis, and a combination of queuing disciplines with FSN analysis were proposed, resulting in seven improvement scenarios. The simulation results indicated that Queuing Approach 1, which prioritizes six-wheeled vehicles in combination with FSN analysis, was the most appropriate method for process improvement. This approach enhances process efficiency compared with the situation prior to the improvement. The average waiting time for all vehicle types was reduced, and the total number of vehicles exiting the system within the specified monthly period increased from 582 to 658. This increase in vehicle throughput enhances the opportunity to sell more products, resulting in a profit of 14,063,037.11 Baht per month for the case study company.

Keywords: FSN analysis, Queuing discipline, Simulation

1) บทนำ

ปัจจุบันตลาดเหล็กโลกมีแนวโน้มปรับราคาลดลงอย่างต่อเนื่อง ทำให้ราคาเหล็กในประเทศไทยต้องเผชิญกับแรงกดดันจากการนำเข้าเหล็กจากต่างประเทศ ส่งผลกระทบต่อปริมาณการผลิตและรายได้ของผู้ค้าเหล็กภายในประเทศ ความสามารถในการทำกำไรจึงกลายเป็นความท้าทายที่เพิ่มขึ้น Kasikorn Research Center [1] ดังนั้นเจ้าของกิจการค้าเหล็กในประเทศจำเป็นต้องปรับปรุงการดำเนินงานเพื่อเพิ่มความสามารถในการแข่งขัน การจัดการคลังสินค้าเป็นสิ่งสำคัญประการหนึ่งในการเพิ่มประสิทธิภาพกระบวนการจัดจำหน่าย ทั้งนี้หากการจัดการคลังสินค้าไม่เป็นระบบ เช่น การจัดวางสินค้าไม่เป็นระเบียบ จะก่อให้เกิดความล่าช้าในการค้นหาสินค้าส่งผลให้การจัดส่งล่าช้า ลดทอนประสิทธิภาพการขาย ด้วยเหตุนี้การจัดการคลังสินค้าจึงเป็นประเด็นที่ควรให้ความสำคัญ และปรับปรุงให้สามารถตอบสนองความต้องการของลูกค้าได้อย่างมีประสิทธิภาพ เป็นการเพิ่มความสามารถในการแข่งขันของธุรกิจ

บริษัทกรณีศึกษา เป็นผู้ผลิตและจำหน่ายผลิตภัณฑ์เหล็กขนาดใหญ่ในประเทศไทย มีผลิตภัณฑ์เหล็กหลากหลายประเภท ถูกจัดเก็บไว้ในคลังขนาดใหญ่เพื่อรอการจัดจำหน่าย โดยรถขนส่งที่เข้ามารับสินค้าและจัดส่งแก่ลูกค้าแบ่งออกเป็น 4 ประเภท ได้แก่ 1) รถสิบล้อ 2) รถเทรลเลอร์ 3) รถหกล้อ 4) รถพ่วงแม่-ลูก ลูกค้าสามารถเข้ามาสั่งซื้อผลิตภัณฑ์ได้ตามต้องการ โดยกระบวนการเริ่มจากการรับคำสั่งซื้อจากลูกค้า และพนักงานจะจัดเตรียมสินค้าในพื้นที่จัดเก็บสินค้า เรียกว่า เฟส (phase) แบ่งออกเป็น 5 เฟส แต่ละเฟสประกอบไปด้วยสินค้า 6 กอง สำหรับวางสินค้าประเภทต่าง ๆ ต่อมารถขนส่งจะเคลื่อนที่ไปขึ้นสินค้าแบบสุ่มยังเฟสที่ว่างและมีประเภทสินค้าตามความต้องการของลูกค้า พนักงานจะลำเลียงสินค้ามายังจุดเตรียมเพื่อรอขึ้นสินค้าขึ้นรถขนส่ง หลังจากนั้นจะนำสินค้าขึ้นรถขนส่งและจัดวางตามรูปแบบที่เหมาะสมสำหรับรถแต่ละประเภท

ปัจจุบันบริษัทกรณีศึกษาพบปัญหาในคลังสินค้า คือ กระบวนการขึ้นสินค้าเพื่อจัดส่งใช้เวลานาน และมีปริมาณรถขนส่งในแถวคอยรอขึ้นสินค้าจำนวนมากในช่วงเวลา 13.00–17.00 น. ส่งผลกระทบต่อคลังสินค้านี้ดังกล่าวเป็นจุดคอขวด จากการวิเคราะห์สาเหตุเบื้องต้น พบว่าเกิดจากการจัดวางสินค้าที่ไม่ระบุตำแหน่ง และไม่มีกระบวนการจัดการแถวคอยที่เหมาะสม ส่งผลให้เวลา

การลำเลียงสินค้าขึ้นรถขนส่งนาน และทำให้เวลารอคอยของรถขนส่งเพิ่มขึ้น ตามลำดับ

ผู้วิจัยได้ตระหนักถึงปัญหาความล่าช้าในกระบวนการขึ้นสินค้าดังกล่าว งานวิจัยนี้จึงจะประยุกต์ใช้การจำลองสถานการณ์ร่วมกับทฤษฎีการจัดการคลังสินค้า และการจัดการแถวคอยเข้ามาแก้ปัญหามานี้ โดยเริ่มจากการรวบรวมข้อมูล เพื่อศึกษารูปแบบกระบวนการที่ยังไม่เหมาะสม วิเคราะห์ปัญหา และหาสาเหตุของปัญหา เมื่อทราบถึงสาเหตุที่แท้จริงแล้วจึงทำการปรับปรุงกระบวนการผ่านการจำลองสถานการณ์ เพื่อเสนอแนวทางการปรับปรุงการจัดวางผลิตภัณฑ์ภายในคลังสินค้า และจัดการแถวคอย ทำให้ระยะเวลาการลำเลียงสินค้า และเวลารอคอยของรถขนส่งลดลง ตามลำดับ ส่งผลทำให้รอบการจัดส่งเพิ่มมากขึ้น

2) ทฤษฎี และงานวิจัยที่เกี่ยวข้อง

2.1) ทฤษฎีที่เกี่ยวข้อง

การจำลองสถานการณ์เป็นการเลียนแบบการดำเนินงานของกระบวนการหรือระบบจริง เพื่อศึกษาพฤติกรรมผ่านการพัฒนาแบบจำลอง [2] และใช้หาแนวทางเพื่อพัฒนาระบบงานให้มีประสิทธิภาพมากขึ้น [3] เหมาะสำหรับการจำลองระบบงานที่ซับซ้อนและมีผลกระทบจากการเปลี่ยนแปลงเชิงข้อมูล องค์กร และสิ่งแวดล้อม [4]

งานวิจัยนี้ใช้การพิจารณาระเบียบแถวคอย (queue discipline) เพื่อกำหนดลำดับความสำคัญของลูกค้า [4] เพื่อเข้ารับบริการภายในคลังสินค้า ในการปรับปรุงแถวคอยของรถขนส่ง และใช้การวิเคราะห์สินค้าคงคลังแบบเอฟเอสเอ็น (FSN analysis) ซึ่งเป็นการแบ่งกลุ่มสินค้าที่อยู่บนพื้นฐานความถี่การเคลื่อนไหวของวัสดุ โดยแบ่งเป็นสินค้าเคลื่อนไหวเร็ว สินค้าเคลื่อนไหวช้า และสินค้าไม่เคลื่อนไหว [5] นำไปสู่การจัดสรรพื้นที่ที่เหมาะสมสำหรับสินค้าขายได้มาก ปานกลาง และน้อย เพื่อเพิ่มประสิทธิภาพกระบวนการลำเลียงสินค้าภายในคลังสินค้าของบริษัทกรณีศึกษา

ในงานวิจัยจะทดสอบความเสมือนจริง (validation) ของแบบจำลองทางคอมพิวเตอร์ จากการทดสอบสมมติฐานทางสถิติแบบ t-test ของค่าเฉลี่ย 2 กลุ่มประชากร [6] โดยเปรียบเทียบดัชนีชี้วัดจากระบบงานจริงกับแบบจำลองทางคอมพิวเตอร์ภายใต้สมมติฐานหลัก (H_0) และสมมติฐานทางเลือก (H_1) ดังนี้

$$H_0 : \mu_1 = \mu_2$$

$$H_1 : \mu_1 \neq \mu_2$$

โดยที่ μ_1 และ μ_2 คือ ค่าเฉลี่ยข้อมูลจากระบบจริงและค่าเฉลี่ยข้อมูลจากแบบจำลอง ตามลำดับ หากไม่สามารถปฏิเสธสมมติฐานหลักได้ ทำให้มั่นใจได้ว่าแบบจำลองทางคอมพิวเตอร์สามารถใช้เป็นตัวแทนของระบบงานจริงได้

2.2) งานวิจัยที่เกี่ยวข้อง

ในอดีตมีผู้วิจัยนำการจำลองสถานการณ์ การแบ่งกลุ่มของคลัง การจัดการแถวคอยไปใช้ปรับปรุงงานคลังสินค้าและระบบงานต่าง ๆ ได้แก่ Zhu *et al.* [7] ปรับปรุงประสิทธิภาพงานโลจิสติกส์ห่วงโซ่ความเย็นของสินค้าผลไม้และผักผ่านการจำลองสถานการณ์ โดยวิเคราะห์หาคอขวดและทรัพยากรที่ไม่ได้ใช้ในการทำงาน Millstein *et al.* [8] เพิ่มประสิทธิภาพการตัดสินใจจัดกลุ่มสินค้าคงคลัง ABC จากโมเดลหาค่าที่เหมาะสม เพื่อตัดสินใจการจัดกลุ่มของคลังที่เหมาะสม Panyapanich and Sirisophonilp [9] ได้สร้างแบบจำลองสถานการณ์สำหรับวิเคราะห์กระบวนการหยิบขึ้นส่วนอิเล็กทรอนิกส์ ปรับเปลี่ยนวิธีการหยิบสินค้า ปรับตำแหน่งจัดเก็บสินค้าและพื้นที่รวบรวมสินค้า ทำให้ระยะเวลารวมและระยะทางในการเดินหยิบสินค้าลดลง การปรับปรุงแผนผังคลังสินค้าสำเร็จรูปและวิเคราะห์ ABC ตามความถี่ในการเคลื่อนไหวของสินค้า เพื่อลดระยะทางรวมของการเคลื่อนไหวของสินค้าถูกทำโดย Wongsamajan [10] การประยุกต์ใช้เทคนิคการจำลองสถานการณ์เพื่อพัฒนารูปแบบการจัดวางสินค้า ทำให้เวลาในการขนย้ายบล็อกแก้วลดลงถูกศึกษาโดย Bunterngrachit [11] Chanyaem and Akarapichet [12] ใช้การวิเคราะห์ ABC จัดหมวดหมู่ของไหลจากความถี่และความถี่ และใช้หลัก ECRS ปรับปรุงกิจกรรมในกระบวนการทำงาน ทำให้ระยะทางการเบิกสินค้าและระยะเวลาการค้นหาลินคาลลดลง Anurattananon *et al.* [13] ปรับตำแหน่งการจัดวางในคลังจัดเก็บเครื่องดื่ม โดยการวิเคราะห์ ABC จากความถี่การซื้อสินค้าและน้ำหนักสินค้า สามารถลดเวลาและต้นทุนภายในคลังสินค้าได้ Petchan and Karoonsoontawong [14] พัฒนาแบบจำลองสถานการณ์เพื่อปรับปรุงระบบแถวคอยบนสถานีรถไฟ ทำให้ผลรวมเวลาในแถวคอย ผลรวมความยาวแถวคอยลดลง ในจุดวิกฤติของการให้บริการ การปรับปรุงกระบวนการไหลสินค้าของธุรกิจวัสดุทดแทนไม้ ถูกศึกษาโดย Kraiwong and Setamanit [15] โดยใช้ผังสายธารคุณค่าและการจำลองสถานการณ์ เพื่อลดระยะเวลาที่ไม่ก่อให้เกิดคุณค่า พบว่าเวลาเฉลี่ยของการเข้ามารับสินค้าและระยะเวลารอคอยเฉลี่ยเพื่อไหลสินค้าประเภท

รถพื้นเรียบและตู้คอนเทนเนอร์ลดลง Ganbold *et al.* [16] วิจัยเพื่อหาค่าเหมาะที่สุดบนพื้นฐานการจำลองสถานการณ์ สำหรับการมอบหมายงานพนักงานคลังสินค้า สามารถปรับปรุงระดับบริการคลังสินค้าได้ Jirapatsil [17] ศึกษาเพื่อปรับปรุงตำแหน่งการวางอาหารสดและกระบวนการขายออกสำหรับคลังสินค้าในธุรกิจพาณิชย์อิเล็กทรอนิกส์ โดยใช้การวิเคราะห์หาค่าตลาดทฤษฎีความสัมพันธ์ของสินค้า รวมถึงใช้หลักการสินค้า สามารถลดระยะทางและระยะเวลาในคลังสินค้า Chaimanee and Kaewcharoensithong [18] ทำการจำลองสถานการณ์เพื่อลดข้อขัดข้องระหว่างกระบวนการแปรรูปข้าวเปลือก ทำให้ข้อขัดข้องโดยเฉลี่ยในกระบวนการลดความชื้น และกระบวนการแปรรูปข้าวลดลง

ต่อมา Preedawat *et al.* [19] ใช้การจำลองสถานการณ์ร่วมกับการวิเคราะห์ FSN เพื่อปรับปรุงประสิทธิภาพการหยิบสินค้า ผลการวิจัยพบว่าระยะเวลาเฉลี่ยต่อคำสั่งซื้อ ระยะทางรวมเฉลี่ยต่อวัน ระยะเวลาเฉลี่ยของคำสั่งซื้อลดลง และทำให้ต้นทุนการจัดส่งรอบพิเศษลดลง Phimkan [20] ใช้เทคนิคการวิเคราะห์ ABC-FSN เพื่อจัดลำดับความสำคัญของชิ้นส่วนซ่อมบำรุงและใช้วิธี Silver-Meal เพื่อจัดการการสั่งซื้อ สามารถลดต้นทุนการจัดการสินค้าคงคลังได้ การปรับปรุงกระบวนการจัดเก็บสินค้าคงคลังโดยการวิเคราะห์ ABC แบ่งกลุ่มตามปริมาณการเคลื่อนไหว ทำให้ระยะทาง และต้นทุนด้านแรงงานลดลง อีกทั้งยังเพิ่มความแม่นยำในการจัดเก็บปรากฏใน Likitkam [21] Nigischer *et al.* [22] สร้างแบบจำลองสถานการณ์ในพื้นที่การทำงานของเครื่อง CNC เริ่มจากแบบจำลองรายละเอียดต่ำ รายละเอียดปานกลาง และรายละเอียดสูงผ่าน FlexSim ผลการวิจัย คือ การจำลองสถานการณ์ที่มีความแม่นยำ Khotyota *et al.* [23] วิจัยเพื่อลดความผิดพลาด และลดระยะเวลาการหยิบสินค้า โดยใช้การวิเคราะห์ ABC และใช้แผนภูมิกระบวนการไหลวิเคราะห์ขั้นตอนการหยิบสินค้า และแก้ไขด้วยหลักการ ECRS ทำให้ความผิดพลาดและระยะเวลาในการหยิบสินค้าลดลง Tuammee *et al.* [24] ปรับปรุงตำแหน่งและการจัดวางวัตถุดิบและสินค้าในคลังสินค้า โดยการวิเคราะห์ ABC สำหรับจัดหมวดหมู่วัตถุดิบและสินค้า ร่วมกับทฤษฎีควบคุมการมองเห็น ช่วยลดระยะเวลาในการค้นหาวัตุดิบและสินค้าได้ Sringean *et al.* [25] วิจัยเพื่อแก้ปัญหาความไม่สอดคล้องระหว่างปริมาณของคลังและปริมาณความต้องการของหน่วยงานราชการ จากการวิเคราะห์ ABC และ FSN สำหรับแบ่งกลุ่มพัสดุตามมูลค่าและอัตราการหมุนเวียน ผลการวิจัยพบว่า

สามารถลดการขาดแคลนและภาวะของสินค้าอีกทั้งยังสามารถทำให้ต้นทุนวัสดุคงคลังลดลง Chatsatayu *et al.* [26] ปรับปรุงตำแหน่งการจัดเก็บบรรจุภัณฑ์โดยการวิเคราะห์ ABC และ FSN และจำลองการหยิบสินค้า พบว่าระยะเวลาเฉลี่ยในการหยิบสินค้าลดลง Ardiansyah *et al.* [27] ใช้การจำลองระบบเหตุการณ์ไม่ต่อเนื่องร่วมกับการวิเคราะห์ ABC แบบสองปัจจัย ปรับตำแหน่งการจัดวางสินค้าใหม่ พบว่าเวลาการหยิบสินค้าของผู้ปฏิบัติงานลดลง Lumsakul *et al.* [28] ปรับปรุงระบบโลจิสติกส์เข้าในโรงงานมะพร้าวผ่านการจำลองสถานการณ์ งานวิจัยพบว่าสามารถเพิ่มผลผลิต ส่งผลให้งานระหว่างกระบวนการ และเวลาในกระบวนการลดลง

จากการทบทวนงานวิจัยในอดีตพบว่ามีการใช้การจัดการระบบแถวคอยหรือการจัดการคลังสินค้าโดยการแบ่งกลุ่มของคงคลังเข้าไปแก้ปัญหาในระบบงานต่าง ๆ สำหรับงานวิจัยนี้ ระบบงานเกิดปัญหาด้านเวลารอคอยของรถขนส่งนานก่อนจะเข้าไปรับสินค้าภายในคลัง และการลำเลียงสินค้าเพื่อนำขึ้นรถขนส่งภายในคลังใช้เวลานาน ซึ่งเป็นการดำเนินงานที่ต่อเนื่องกัน ในงานวิจัยจึงจะนำการพิจารณาระเบียบแถวคอยไปใช้ร่วมกับการปรับปรุงการจัดวางสินค้าภายในคลังผ่านการจำลองสถานการณ์ เพื่อปรับปรุงกระบวนการจัดส่งสินค้าของบริษัทกรณีศึกษา วัตถุประสงค์เพื่อลดเวลารอคอยของรถขนส่งและระยะเวลาการลำเลียงสินค้าส่งผลให้ผลประโยชน์ของบริษัทเพิ่มขึ้น

3) วิธีดำเนินการวิจัย

3.1) กระบวนการขึ้นสินค้า

งานวิจัยเริ่มจากศึกษากระบวนการขึ้นสินค้าปัจจุบัน สำหรับสินค้าประเภทเหล็ก โดยผลิตภัณฑ์จะถูกเก็บไว้ในคลังขนาดใหญ่เพื่อรอการจัดจำหน่ายสู่ลูกค้า ทั้งนี้รถขนส่งจะเข้าไปรับสินค้าภายในคลัง การดำเนินงานแสดงดังแผนภูมิการไหลในรูปที่ 1

จากแผนภูมิการไหลในรูปที่ 1 สามารถอธิบายรายละเอียดการดำเนินงานในแต่ละขั้นตอน ดังนี้

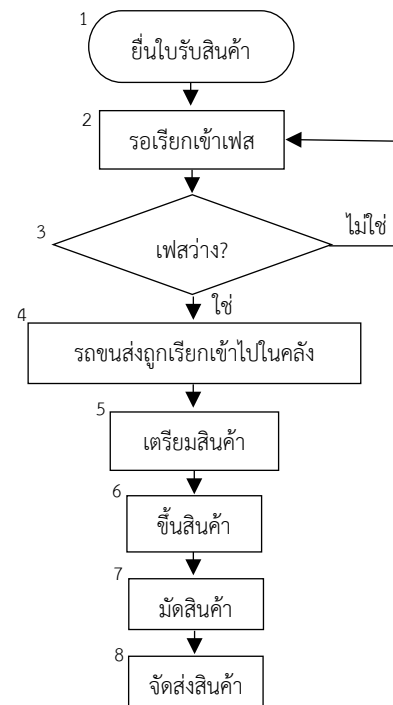
3.1.1) ขั้นตอนที่ 1: รถขนส่งมาถึงคลังสินค้า ยื่นใบรับสินค้าให้แก่พนักงาน

3.1.2) ขั้นตอนที่ 2: รอพนักงานเรียกรถขนส่งเข้าเฟส

3.1.3) ขั้นตอนที่ 3: ตรวจสอบว่ามีเฟสว่างหรือไม่

3.1.4) ขั้นตอนที่ 4: รถขนส่งถูกเรียกเข้าไปในคลังเมื่อมีเฟสว่าง

3.1.5) ขั้นตอนที่ 5: พนักงานใช้เครนลำเลียงสินค้ามาที่จุดเตรียมขึ้นสินค้าตามความต้องการของลูกค้า



รูปที่ 1 : แผนภูมิการไหลแสดงขั้นตอนการทำงาน

3.1.6) ขั้นตอนที่ 6: พนักงานใช้เครนนำสินค้าขึ้นรถขนส่งและจัดวางสินค้าบนรถตามรูปแบบที่เหมาะสมกับรถแต่ละประเภท

3.1.7) ขั้นตอนที่ 7: พนักงานมัดสินค้าให้แน่นเพื่อป้องกันความเสียหายของสินค้าระหว่างการขนส่ง

3.1.8) ขั้นตอนที่ 8: จัดส่งสินค้าไปสู่ลูกค้า

จากกระบวนการทำงานที่กล่าวมาข้างต้น พบว่าปัญหาที่เกิดขึ้นในกระบวนการขึ้นสินค้า แบ่งออกเป็น 2 กรณี ได้แก่

กรณีที่ 1 เวลารอคอยของรถขนส่งในแถวคอยยาวนาน

กรณีที่ 2 การลำเลียงสินค้ามาที่จุดเตรียมใช้เวลานาน

ปัญหาทั้ง 2 กรณี ดังกล่าว ส่งผลทำให้เวลารอคอยของรถขนส่งยาวนาน และทำให้จำนวนรถขนส่งที่ออกจากระบบภายในช่วงเวลากรณีศึกษามีปริมาณน้อย

จากปัญหาที่กล่าวข้างต้น งานวิจัยนี้จะพัฒนากระบวนการตัดสินใจเพื่อจัดลำดับความสำคัญของประเภทรถขนส่งภายในแถวคอย และการระบุตำแหน่งการจัดวางสินค้าภายในคลัง เพื่อลดเวลารอคอย และทำให้การลำเลียงสินค้าเร็วขึ้น ส่งผลให้จำนวนรถที่ออกจากระบบเพิ่มขึ้น

3.2) การเก็บข้อมูลและการสร้างแบบจำลองสถานการณ์

เมื่อศึกษากระบวนการขึ้นสินค้าแล้ว จะเก็บข้อมูลเพื่อนำมาสร้างแบบจำลองทางคอมพิวเตอร์ผ่านโปรแกรม FlexSim โดยรวบรวมข้อมูลย้อนหลัง 2 เดือน โดยจะใช้ข้อมูลในช่วงเวลา 13.00–17.00 น. ซึ่งเป็นช่วงเวลาที่ปริมาณรถขนส่งมารับสินค้าจำนวนมาก ทำให้เกิดการรอคอยของรถขนส่งและเกิดปัญหาความล่าช้าในการลำเลียงสินค้าภายในคลัง แตกต่างจากช่วงเวลาอื่น ๆ ซึ่งมีจำนวนรถเข้ามารับสินค้าจำนวนน้อยทำให้เกิดปัญหาความล่าช้าในการดำเนินงาน ข้อมูลที่ต้องใช้เพื่อสร้างแบบจำลองระบบงานภายในคลังมีดังนี้

3.2.1) ประเภทของรถขนส่ง (Vehicle Type): แบ่งเป็น 4 ประเภท ได้แก่ รถสิบล้อ รถเทรลเลอร์ รถหกล้อ และรถพ่วงแม่-ลูก โดยกำหนดสัญลักษณ์สำหรับรถแต่ละประเภท คือ A, B, C และ D ตามลำดับ

3.2.2) ช่วงเวลาห่างของรถแต่ละประเภท (Interarrival Time): เพื่อแสดงระยะห่างของเวลาระหว่างการเข้ามาของรถขนส่ง สะท้อนถึงความถี่ในการเข้ามาของรถแต่ละประเภท

3.2.3) เวลาเตรียมสินค้า (Setup Time): แสดงถึงเวลาที่พนักงานเริ่มลำเลียงสินค้ามาที่จุดเตรียมสินค้าตามคำสั่งซื้อของลูกค้าจนเสร็จสิ้น

3.2.4) เวลาขึ้นสินค้า (Processing Time): แสดงถึงเวลาที่พนักงานเริ่มนำสินค้าขึ้นรถขนส่งจนครบจำนวนตามคำสั่งซื้อ

3.2.5) ประเภทของสินค้า (Product Type): แสดงถึงสัดส่วนของสินค้าที่ขายได้มาก ปานกลาง และน้อย ในแต่ละเฟส

ข้อมูลข้างต้นจะถูกใช้เป็นพารามิเตอร์นำเข้า เพื่อสร้างแบบจำลองทางคอมพิวเตอร์ โดยนำข้อมูลไปสร้างการแจกแจงความน่าจะเป็นที่เหมาะสมแสดงพฤติกรรมการทำงานปัจจุบันภายในคลังสินค้า โดยพิจารณาจากค่าคะแนนความสัมพันธ์ (Relative Score) หากมีค่าใกล้เคียง 100 เปอร์เซนต์ มากที่สุด แสดงว่าข้อมูลเหมาะสมกับการแจกแจงประเภทนั้นๆ ยกตัวอย่างการแจกแจงความน่าจะเป็นช่วงเวลาห่างของรถ เวลาเตรียมสินค้า และเวลาขึ้นสินค้าของรถประเภท D แสดงดังรูปที่ 2–4 และสำหรับรถทุกประเภท สรุปดังตารางที่ 1

หลังจากนั้นนำข้อมูลการแจกแจงความน่าจะเป็นดังกล่าวไปสร้างแบบจำลองทางคอมพิวเตอร์ แสดงดังรูปที่ 5 ซึ่งสามารถอธิบายส่วนประกอบและลักษณะการดำเนินงานภายในคลังสินค้าสรุปได้ดังตารางที่ 2

Relative Evaluation of Candidate Models

Model	Relative Score	Parameters
1 - Beta	98.86	Lower endpoint
		Upper endpoint
		Shape #1
		Shape #2
2 - Johnson SB	96.59	Lower endpoint
		Upper endpoint
		Shape #1
		Shape #2
3 - Gamma	71.59	Location
		Scale
		Shape

23 models are defined with scores between 2.27 and 98.86

Absolute Evaluation of Model 1 - Beta

Evaluation: Good

Suggestion: Additional evaluations using Comparisons Tab might be informative.

See Help for more information.

Additional Information about Model 1 - Beta

"Error" in the model mean

relative to the sample mean 7.54316 = 5.07%

รูปที่ 2 : การแจกแจงความน่าจะเป็นช่วงเวลาห่างของรถประเภท D

Relative Evaluation of Candidate Models

Model	Relative Score	Parameters
1 - Pearson Type VI(E)	96.43	Location
		Scale
		Shape #1
		Shape #2
2 - Johnson SB	90.48	Lower endpoint
		Upper endpoint
		Shape #1
		Shape #2
3 - Inverse Gaussian(E)	84.52	Location
		Scale
		Shape

22 models are defined with scores between 0.00 and 96.43

Absolute Evaluation of Model 1 - Pearson Type VII(E)

Evaluation: Good

Suggestion: Additional evaluations using Comparisons Tab might be informative.

See Help for more information.

Additional Information about Model 1 - Pearson Type VII(E)

"Error" in the model mean

relative to the sample mean -0.27418 = 1.02%

รูปที่ 3 : การแจกแจงความน่าจะเป็นเวลาเตรียมสินค้าของรถประเภท D

Relative Evaluation of Candidate Models

Model	Relative Score	Parameters
1 - Beta	100.00	Lower endpoint
		Upper endpoint
		Shape #1
		Shape #2
2 - Johnson SB	96.30	Lower endpoint
		Upper endpoint
		Shape #1
		Shape #2
3 - Weibull	76.85	Location
		Scale
		Shape

28 models are defined with scores between 0.00 and 100.00

Absolute Evaluation of Model 1 - Beta

Evaluation: Good

Suggestion: Additional evaluations using Comparisons Tab might be informative.

See Help for more information.

Additional Information about Model 1 - Beta

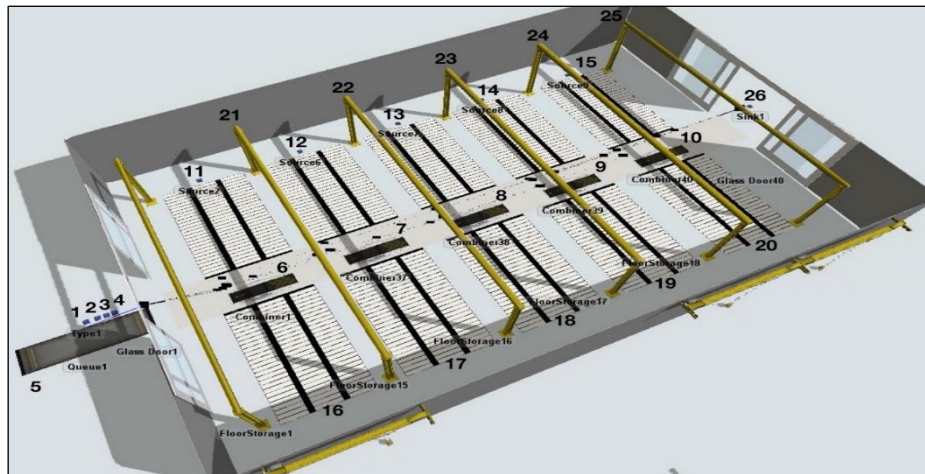
"Error" in the model mean

relative to the sample mean 0.21594 = 0.55%

รูปที่ 4 : การแจกแจงความน่าจะเป็นเวลาขึ้นสินค้าของรถประเภท D

ตารางที่ 1 : การแจกแจงความน่าจะเป็นสำหรับช่วงเวลาห่าง เวลาการเตรียมสินค้า และเวลาการขึ้นสินค้าของรถแต่ละประเภท

ข้อมูล	ประเภทรถ	ประเภทการแจกแจงที่เหมาะสม
ช่วงเวลาห่างของรถ (นาที) (Interarrival Time)	รถสิบล้อ (A)	Beta(3.08651, 191.55683, 0.89684, 7.30459)
	รถเทรลเลอร์ (B)	Beta(0.27297, 301.35409, 0.83920, 8.07690)
	รถหกล้อ (C)	Gamma(0.0, 32.75953, 1.74243)
	รถพ่วงแม่-ลูก (D)	Beta(15.36466, 283.46481, 0.45593, 0.51606)
เวลาเตรียมสินค้า (นาที) (Setup Time)	รถสิบล้อ (A)	Pearson6(3.98631, 85.87169, 0.82770, 4.15546)
	รถเทรลเลอร์ (B)	Pearson6(3.99362, 88.54528, 0.83983, 3.95378)
	รถหกล้อ (C)	Weibull(3.99100, 15.28981, 0.68152)
	รถพ่วงแม่-ลูก (D)	Pearson6(3.77237, 60.86783, 1.05071, 3.74109)
เวลาขึ้นสินค้า (นาที) (Processing Time)	รถสิบล้อ (A)	Pearson6(0.00000, 8.39141, 4.02246, 2.53642)
	รถเทรลเลอร์ (B)	Pearson6(0.00000, 22.81452, 3.31300, 3.40418)
	รถหกล้อ (C)	Lognormal2(3.31800, 8.60585, 0.98274)
	รถพ่วงแม่-ลูก (D)	Beta(11.87651, 93.49780, 0.62687, 1.26957)



รูปที่ 5: แบบจำลองทางคอมพิวเตอร์ผ่านโปรแกรม FlexSim Simulation

ตารางที่ 2 : รายละเอียดของแบบจำลอง

หมายเลข	คำอธิบาย
1-4	สร้างรถแต่ละประเภทเข้ามาในคลังสินค้า
5	แถวคอยเพื่อรอเรียกรถเข้าเฟส
6-10	พื้นที่สำหรับการนำสินค้าขึ้นรถขนส่ง
11-15	สร้างสินค้าประเภท F, S, N
16-20	เฟสสำหรับเก็บสินค้า
21-25	เครนสำหรับลำเลียงสินค้าขึ้นรถขนส่ง
26	รถขนส่งออกจากระบบ

3.3) การตรวจสอบความถูกต้องและการทวนสอบความเหมือนจริงของแบบจำลองทางคอมพิวเตอร์

การตรวจสอบความถูกต้องเพื่อให้มั่นใจได้ว่ากระบวนการทำงานภายในคลังสินค้ามีลำดับขั้นตอนที่ถูกต้อง โดยสังเกตการ

ไหลของรถขนส่ง (entity) ที่เคลื่อนที่เข้าไปยังเฟสที่ว่าง มีการเตรียมสินค้าและการขึ้นสินค้าตามที่กำหนด จากนั้นจึงทวนสอบความเหมือนจริงเพื่อเป็นการยืนยันว่าแบบจำลองทางคอมพิวเตอร์มีการประมวลผลแสดงพฤติกรรมที่แม่นยำ [29] โดยเปรียบเทียบผลลัพธ์จากแบบจำลองผ่านการประมวลผล 35 รอบทำซ้ำ และข้อมูลที่เก็บจากระบบการจริง ในช่วงเวลา 13.00–17.00 น. โดยดัชนีชี้วัดการทวนสอบ ได้แก่ เวลารอคอยและจำนวนรถขนส่งแต่ละประเภทที่ออกจากระบบ โดยการทดสอบสมมติฐานค่าเฉลี่ยของ 2 กลุ่มประชากร จากวิธี t-test ที่ช่วงความเชื่อมั่น 95 เปอร์เซ็นต์ ผ่านโปรแกรม Minitab ผลการทวนสอบความเหมือนจริงแสดงดังตารางที่ 3

ตารางที่ 3 : การทวนสอบความเหมือนจริงของแบบจำลองทางคอมพิวเตอร์

ประเภทของรถ	ดัชนีชี้วัด	ระบบจริง	แบบจำลอง	ผลการทดสอบ	
				P-value	สรุป
รถสิบล้อ (A)	เวลาการรอคอยของรถ (นาทิจ)	14.12 ± 2.75	15.18 ± 1.02	0.586	Do not reject H_0
	จำนวนรถที่ออกจากระบบ (คัน)	7.00 ± 1.62	8.11 ± 2.03	0.220	Do not reject H_0
รถเทรลเลอร์ (B)	เวลาการรอคอยของรถ (นาทิจ)	21.44 ± 5.71	20.91 ± 3.60	0.872	Do not reject H_0
	จำนวนรถที่ออกจากระบบ (คัน)	5.83 ± 0.77	6.37 ± 0.83	0.397	Do not reject H_0
รถหกล้อ (C)	เวลาการรอคอยของรถ (นาทิจ)	15.11 ± 5.16	15.64 ± 5.33	0.898	Do not reject H_0
	จำนวนรถที่ออกจากระบบ (คัน)	3.07 ± 0.80	3.31 ± 0.61	0.619	Do not reject H_0
รถพ่วงแม่-ลูก (D)	เวลาการรอคอยของรถ (นาทิจ)	19.30 ± 11.35	20.81 ± 8.51	0.853	Do not reject H_0
	จำนวนรถที่ออกจากระบบ (คัน)	1.21 ± 0.25	1.62 ± 0.26	0.074	Do not reject H_0

จากตารางที่ 3 พบว่าการประมาณค่าแบบช่วงจากระบบจริง และแบบจำลองทางคอมพิวเตอร์มีช่วงที่ซ้อนทับกัน และผลการทดสอบสมมติฐานยังแสดงให้เห็นว่าดัชนีชี้วัดจากระบบปัจจุบัน และแบบจำลองไม่แตกต่างกันอย่างมีนัยสำคัญ ($P\text{-value} > 0.05$) ทำให้มั่นใจได้ว่าแบบจำลองที่สร้างขึ้นมีความเหมือนจริงทางสถิติสามารถใช้เป็นตัวแทนของกระบวนการขึ้นสินค้าปัจจุบันได้

หลังจากนั้นจะใช้ Why-Why Analysis วิเคราะห์ปัญหา แสดงดังตารางที่ 4 นำไปสู่แนวทางการแก้ไข

ตารางที่ 4 : Why-Why Analysis

what	why	why
เวลารอคอยของรถก่อนเข้าคลังสินค้านาน	ใช้กฎมาก่อนเข้าก่อน (พิจารณาทุกประเภทร่วมกัน)	ไม่มีการจัดลำดับความสำคัญของรถเข้าคลังสินค้าอย่างเหมาะสม
การล่าเลยสินค้าจากพื้นที่จัดเก็บใช้เวลานาน	การจัดเก็บสินค้าไม่เป็นระบบทำให้ใช้เวลาล่าเลยสินค้านาน	-

จาก Why-Why Analysis นำไปสู่แนวทางการปรับปรุงกระบวนการทำงาน รายละเอียดแสดงในหัวข้อถัดไป

3.4) การจัดการแถวคอย

การจัดการแถวคอยเริ่มจากกำหนดเกณฑ์การตัดสินใจ โดยพิจารณาช่วงห่างของการมาถึงและระยะเวลาที่รถขนส่งแต่ละประเภทอยู่ในคลัง (เวลาเตรียมงานรวมกับเวลาขึ้นงาน) แสดงดังตารางที่ 5

ตารางที่ 5 : เกณฑ์การพิจารณาการจัดแถวคอย

ประเภทรถ	ช่วงห่างการมาถึง (นาทิจ)	เวลาที่รถอยู่ในคลัง (นาทิจ)
A	[1,30] : สูง	[50,100] : กลาง
B	[1,30] : สูง	[50,100] : กลาง
C	[31,50] : กลาง	[11,49] : ต่ำ
D	[51,120] : ต่ำ	[101,147] : สูง

จากตารางที่ 5 แสดงช่วงห่างการมาถึงของรถขนส่งหากมีค่าน้อยสะท้อนถึงอัตราการเข้ามาของรถขนส่งสูง และระยะเวลาทั้งหมดที่รถขนส่งอยู่ในคลังก่อนออกจากระบบ ยกตัวอย่างรถประเภท D มีช่วงห่างของการมาถึงอยู่ในช่วง 51–120 นาทิจ แสดงให้เห็นว่ารถประเภทนี้มีอัตราการเข้ามาต่ำที่สุด เวลาที่รถอยู่ในคลังสำหรับรถประเภท D อยู่ในช่วง 101–147 นาทิจ หลังจากนั้นนำช่วงเวลาห่างการมาถึง เวลาที่รถอยู่ในคลัง เป็นเกณฑ์การให้ความสำคัญสำหรับรถบางประเภทในแถวคอยให้สามารถเข้าสู่คลังได้ก่อนรถประเภทอื่น ๆ วัตถุประสงค์เพื่อมุ่งเน้นลดเวลารอคอย โดยแนวทางการจัดการแถวคอยแสดงดังตารางที่ 6 และมีรายละเอียด ดังนี้

3.4.1) การจัดการแถวคอยที่ 1: ให้ความสำคัญกับรถประเภท C เนื่องจากมีช่วงห่างการมาถึงของรถปานกลาง แต่ใช้เวลาในคลังสินค้าสั้นที่สุด จึงมีแนวคิดว่ารรถขนส่งประเภทอื่นๆ จะเกิดการรอคอยต่ำหากรอขึ้นสินค้าตามหลังรถประเภท C แนวคิดนี้เป็นการเร่งระบายรถออกจากระบบ ลดความหนาแน่นของแถวคอย แต่เพื่อป้องกันไม่ให้เกิดที่ช่วงห่างการมาถึงต่ำกว่าเกิดการรอคอย จึงกำหนดให้รถประเภท C เข้าคลังสินค้าได้สูงสุด 4 เฟส ขณะที่รถประเภทอื่นเข้าได้ทุกเฟส ถ้าเฟสนั้น ๆ ว่าง

ตารางที่ 6 : รูปแบบการจัดระเบียบแถวคอยแสดงประเภทรถขนส่งที่เข้าไปรับสินค้าได้ในแต่ละเฟส

เฟส	รูปแบบการจัดแถวคอย (ประเภทรถที่ถูกให้ความสำคัญ)		
	การจัดแถวคอยที่ 1 (ให้ความสำคัญกับ C)	การจัดแถวคอยที่ 2 (ให้ความสำคัญกับ D)	การจัดแถวคอยที่ 3 (ให้ความสำคัญกับ D)
1	A,B,D	A,B,D	A,B,C,D
2	A,B,C,D	A,B,C,D	A,B,C,D
3	A,B,C,D	A,B,C,D	A,B,C,D
4	A,B,C,D	A,B,C,D	A,B,C,D
5	A,B,C,D	A,B,C,D	A,B,C,D

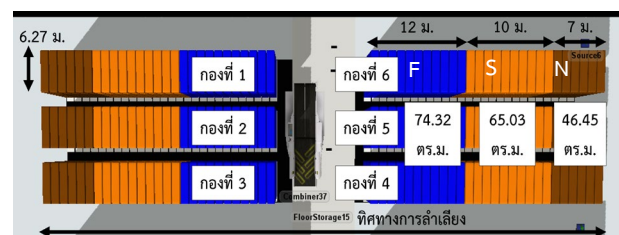
3.4.2) การจัดแถวคอยที่ 2: ให้ความสำคัญกับรถประเภท D เนื่องจากเมื่อวิเคราะห์ช่วงทางการมาถึงพบว่ารถประเภทนี้จะเข้าสู่ระบบ 2-4 คัน (ต่ำที่สุด) และใช้เวลาในคลังยาวนาน โดยมีแนวคิดเพื่อให้รถที่ใช้เวลาในคลังนานไม่เกิดการรอคอย สามารถขึ้นงานได้แล้วเสร็จในช่วง 13.00-17.00 น. เป็นการเพิ่มจำนวนรถที่ออกจากระบบ ทั้งนี้จะไม่เกิดผลกระทบต่อการทำงานของรถขนส่งประเภทอื่นเนื่องจากรถประเภท D อัตราการเข้ามาต่ำที่สุด จึงกำหนดให้สามารถเข้าสู่เฟสที่ว่างได้ก่อนรถประเภทอื่น ๆ ทั้งนี้ยังคงกำหนดให้รถประเภท C เข้าได้สูงสุด 4 เฟส เพื่อรักษาสสมดุลระหว่างการเร่งระบายรถที่ขึ้นสินค้าได้เร็วและทำให้รถที่ใช้เวลาในคลังนานดำเนินงานจนแล้วเสร็จในช่วงเวลาที่กำหนด

3.4.3) การจัดแถวคอยที่ 3: อ้างอิงแนวคิดของการจัดแถวคอยที่ 2 คือให้ความสำคัญกับรถประเภท D แต่กำหนดให้รถทุกประเภทเข้าได้ทุกเฟส แนวทางนี้เพิ่มความยืดหยุ่นในการจัดแถวคอย ลดการว่างของเฟส เป็นการช่วยลดเวลารอคอยของรถทุกประเภท

3.5) การจัดการคลังสินค้าโดยการวิเคราะห์ FSN

การจัดการพื้นที่จัดเก็บสินค้าแบบ FSN มีเป้าหมายเพื่อลดเวลาการลำเลียงสินค้า ส่งผลให้ปริมาณรถออกจากระบบได้มากขึ้น สินค้าประเภท F, S และ N คือ สินค้าที่ขายได้มาก ปานกลาง และน้อย ตามลำดับ โดยสินค้าประเภท F จะถูกจัดวางในพื้นที่จัดเก็บใกล้กับบริเวณที่รถขนส่งเข้ามารับสินค้ามากที่สุด รองลงมาคือสินค้าประเภท S และสินค้าประเภท N จะถูกจัดวางไกลจากพื้นที่ที่จัดส่งมากที่สุด การวิเคราะห์ FSN เริ่มจากการวิเคราะห์สัดส่วนแบ่งตามพื้นที่ของสินค้าแต่ละประเภทในแต่ละเฟส นำไปสู่การหาเวลาการเคลื่อนที่ ยกตัวอย่างการคำนวณในเฟสที่ 2 ดังนี้

- 1) พื้นที่เก็บสินค้าในเฟส = 1,114.80 ตารางเมตร
- 2) พื้นที่เก็บสินค้าแยกตามประเภท F, S, N คือ 445.92, 390.18, 278.70 ตารางเมตร ตามลำดับ
- 3) พื้นที่เก็บสินค้าแยกตามประเภท F, S, N ในแต่ละกอง คือ 74.32, 65.03, 46.45 ตารางเมตร ตามลำดับ แสดงดังรูปที่ 6



รูปที่ 6 : พื้นที่แยกประเภท FSN สำหรับเฟส 2

- 4) เวลาเฉลี่ยการลำเลียงสินค้าจากระยะการเคลื่อนที่ตามทิศทางการลำเลียง (พิจารณาจากพื้นที่จัดเก็บ)

ยกตัวอย่างการคำนวณเวลาเฉลี่ยการลำเลียงสินค้าในเฟสที่ 2 สินค้าประเภท F คำนวณจากค่าเฉลี่ยของระยะทางที่จุดเตรียมสินค้า 3.50 เมตร ระยะไกลสุดของพื้นที่จัดเก็บประมาณ 12 เมตร ระยะกึ่งกลางของพื้นที่จัดเก็บ 6 เมตร จะได้ระยะทางเฉลี่ย 7.17 เมตร และสินค้าประเภท S คำนวณจากค่าเฉลี่ยของระยะทางที่จุดเตรียมสินค้า 3.50 เมตร ระยะจุดเริ่มต้นของพื้นที่จัดเก็บ 12 เมตร ระยะไกลสุดของพื้นที่จัดเก็บประมาณ 22 เมตร ระยะกึ่งกลางของพื้นที่เก็บสินค้า 17 เมตร จะได้ระยะทางเฉลี่ย 13.63 เมตร ทั้งนี้ความเร็วเคลื่อน คือ 1.12 เมตร/วินาที จะได้เวลาเฉลี่ยการเคลื่อนที่เพื่อลำเลียงสินค้าประเภท F และ S เท่ากับ 6.40 วินาที และ 12.17 วินาที ตามลำดับ ผลการคำนวณเวลาการเคลื่อนที่ลำเลียงสินค้าในแต่ละเฟสแสดงดังตารางที่ 7

จากเวลาเฉลี่ยการเคลื่อนที่ จะคำนวณเวลาการลำเลียงสินค้าสำหรับรถแต่ละประเภทตามสมการที่ (1)

$$T_{Crane,(min,max)} = 2 \sum_{i=F}^N n_i t_{x,i} + 2t_z \sum_{i=F}^N n_i + t_h \sum_{i=F}^N n_i + \varepsilon \quad (1)$$

โดยที่ $T_{Crane,(min,max)}$ = เวลาที่ใช้ในการลำเลียงสินค้าน้อยที่สุดหรือมากที่สุด

n_i = จำนวนครั้งในการลำเลียงสินค้าประเภท

$i; i = F, S, N$

$t_{x,i}$ = เวลาเฉลี่ยการเคลื่อนเครื่องเพื่อลำเลียงสินค้าประเภท $i; i = F, S, N$

t_z = เวลาเคลื่อนเครื่องขึ้น-ลง (23 วินาที)

t_h = เวลาคล้องสายพานกับสินค้า (120 วินาที)

ε = เวลาเผื่อ (177.50 วินาที)

ยกตัวอย่างการคำนวณสำหรับรถประเภท D มีการลำเลียงสินค้าแบ่งเป็นเปอร์เซ็นต์โดยเฉลี่ย ได้แก่ $F = 44\%$, $S = 30\%$ และ $N = 26\%$ และมีจำนวนครั้งการลำเลียงน้อยที่สุดและมากที่สุด 5 ครั้งและ 10 ครั้ง ตามลำดับ

การลำเลียง 5 ครั้ง แบ่งเป็นการลำเลียงสินค้าประเภท $F \approx 2$, $S \approx 2$ และ $N \approx 1$ การลำเลียง 10 ครั้ง แบ่งเป็นการลำเลียงสินค้าประเภท $F \approx 4$, $S \approx 3$ และ $N \approx 3$

$$T_{crane,min} = 2 \times [(2 \times 6.40) + (2 \times 12.17) + (1 \times 17.86)] + (2 \times 23 \times 5) + (120 \times 5) + 177.50 = 1117.50 \text{ วินาที (18.63 นาที)}$$

$$T_{crane,max} = 2 \times [(4 \times 6.40) + (3 \times 12.17) + (3 \times 17.86)] + (2 \times 23 \times 10) + (120 \times 10) + 177.50 = 2068.88 \text{ วินาที (34.48 นาที)}$$

จากความเร็วเครื่องที่คำนวณได้ กำหนดเวลาการลำเลียงเป็นการแจกแจงแบบเอกรูปที่มีเวลาการลำเลียงน้อยที่สุด 18.63 นาที และมากที่สุด 34.48 นาที กำหนดเข้าสู่แบบจำลองทางคอมพิวเตอร์ สำหรับรถประเภทอื่น ๆ มีการคำนวณในลักษณะเดียวกัน สรุปได้ตามตารางที่ 8

3.6) การจัดการแถวคอยร่วมกับการวิเคราะห์ FSN

แนวคิดนี้เกิดจากปัญหาวานวิจัยซึ่งแบ่งเป็น 2 กรณี ได้แก่ เวลาในแถวคอยของรถขนส่งนาน และการลำเลียงสินค้าภายในคลังใช้เวลานาน ซึ่งปัญหาที่เกิดขึ้นอยู่ในการดำเนินงานที่ต่อเนื่องกัน เพื่อที่จะแก้ปัญหาทั้ง 2 กรณี ไปพร้อมกัน งานวิจัยนี้จึงนำการปรับปรุงแถวคอยและการวิเคราะห์ FSN มาใช้ร่วมกัน และสามารถสร้างสถานการณ์ปรับปรุงได้อีก 3 แนวทาง ได้แก่ การจัดการแถวคอยที่ 1 - การวิเคราะห์ FSN, การจัดการแถวคอยที่ 2 - การวิเคราะห์ FSN และการจัดการแถวคอยที่ 3 - การวิเคราะห์ FSN

ตารางที่ 7 : ส่วนประกอบนำไปสู่การคำนวณเวลาการลำเลียงสินค้า

เฟส	พื้นที่เก็บสินค้าทั้งหมด (ตร.ม.)	ประเภทสินค้า	พื้นที่เก็บสินค้าแยกตามประเภท (ตร.ม.)	พื้นที่เก็บสินค้าแยกตามประเภท/กอง	ระยะการเคลื่อนเครื่อง (เมตร)	เวลาเฉลี่ยการเคลื่อนเครื่อง ($t_{x,i}$, วินาที)
1	1,000.06	F	422.25	70.38	3.50, 5.50, 11	5.95
		S	200.01	33.34	3.50, 11, 13.50, 16	9.82
		N	377.80	62.97	3.50, 16, 21, 26	14.84
2	1,114.80	F	445.92	74.32	3.50, 6, 12	6.40
		S	390.18	65.03	3.50, 12, 17, 22	12.17
		N	278.70	46.45	3.50, 22, 25.50, 29	17.86
3	1,102.42	F	588.37	98.06	3.50, 8, 16	8.18
		S	433.54	72.26	3.50, 16, 22, 28	15.51
		N	80.51	13.42	3.50, 28, 29, 30	20.20
4	1,114.81	F	421.15	70.19	3.50, 5.50, 11	5.95
		S	377.80	62.97	3.50, 11, 16, 21	11.50
		N	315.86	52.64	3.50, 21, 25, 29	17.52
5	1,114.81	F	501.66	83.61	3.50, 6.50, 13	6.85
		S	365.41	60.90	3.50, 13, 18, 23	12.83
		N	247.74	41.29	3.50, 23, 26.50, 30	18.53

ตารางที่ 8 : การแจกแจงเวลาการลำเลียงสินค้าของแต่ละประเภทหลังวิเคราะห์ FSN

เฟส	ประเภทรถ	จำนวนครั้งในการขึ้นสินค้า/คัน	การแจกแจงข้อมูล
1	A	3-5	uniform(12.28,18.34)
	B	5-8	uniform(18.34,27.66)
	C	2-5	uniform(9.02,18.34)
	D	5-10	uniform(18.34,33.88)
2	A	3-5	uniform(12.47,18.63)
	B	5-8	uniform(18.63,28.14)
	C	2-5	uniform(9.11,18.63)
	D*	5-10	uniform(18.63,34.48)
3	A	3-5	uniform(12.72,19.04)
	B	5-8	uniform(19.04,28.81)
	C	2-5	uniform(9.28,19.04)
	D	5-10	uniform(19.04,35.29)
4	A	3-5	uniform(12.42,18.54)
	B	5-8	uniform(18.54,28.00)
	C	2-5	uniform(9.07,18.54)
	D	5-10	uniform(18.54,34.32)
5	A	3-5	uniform(12.53,18.72)
	B	5-8	uniform(18.72,28.30)
	C	2-5	uniform(9.15,18.72)
	D	5-10	uniform(18.72,34.67)

หมายเหตุ: *ตัวอย่างการคำนวณ

4) ผลการวิจัยและอภิปรายผล

จากกระบวนการจัดแถวคอยและการวิเคราะห์ FSN นำไปสู่สถานการณ์การปรับปรุง 7 สถานการณ์ โดยใช้เวลารอคอยเฉลี่ยของรถแต่ละประเภทในแถวคอย และจำนวนรถที่ออกจากระบบ

โดยเฉลี่ยเป็นดัชนีชี้วัดการปรับปรุงกระบวนการ ผลการจำลองสถานการณ์ แสดงดังตารางที่ 9

จากตารางที่ 9 สามารถวิเคราะห์ผลจากการจำลองสถานการณ์เทียบกับสถานการณ์ก่อนปรับปรุงสรุปตามแนวทางได้ ดังนี้

ตารางที่ 9: เวลารอคอยเฉลี่ยและจำนวนรถที่ออกจากระบบเฉลี่ยจากการจำลองสถานการณ์

ดัชนีชี้วัด	ประเภทรถ	สถานการณ์							
		ก่อนปรับปรุง	การจัดแถวคอยที่ 1	การจัดแถวคอยที่ 2	การจัดแถวคอยที่ 3	วิเคราะห์ FSN	การจัดแถวคอยที่ 1-FSN	การจัดแถวคอยที่ 2-FSN	การจัดแถวคอยที่ 3-FSN
เวลารอคอยเฉลี่ย	A	15.18	16.41	14.22	15.45	11.10	11.03	11.41	12.04
	B	20.91	20.19	18.92	19.60	14.76	16.18	15.91	15.40
	C	15.64	7.28	18.30	15.53	12.81	5.40	12.28	12.88
	D	20.81	17.41	8.06	8.20	12.81	15.02	5.32	6.45
จำนวนรถที่ออกจากระบบเฉลี่ย	A	8.11	8.40	8.54	8.57	9.11	9.03	9.06	9.11
	B	6.37	6.77	6.91	6.49	7.17	7.37	7.09	7.00
	C	3.31	3.59	3.36	3.29	3.38	3.63	3.44	3.38
	D	1.62	1.50	1.76	1.77	1.74	1.89	1.92	1.88

4.1) การจัดแถวคอยที่ 1

4.1.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท B, C และ D ลดลง 0.72 นาที, 8.36 นาที และ 3.40 นาที ตามลำดับ แต่เวลารอคอยของรถประเภท A เพิ่มขึ้น 1.23 นาที

4.1.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B และ C เพิ่มขึ้น 0.29 คัน 0.40 คัน และ 0.28 คัน ตามลำดับ แต่จำนวนรถประเภท D ลดลง 0.12 คัน

4.2) การจัดแถวคอยที่ 2

4.2.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท A, B และ D ลดลง 0.96 นาที, 1.99 นาที และ 12.75 นาที ตามลำดับ แต่เวลารอคอยของรถประเภท C เพิ่มขึ้น 2.66 นาที

4.2.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B, C และ D เพิ่มขึ้น 0.43 คัน, 0.54 คัน, 0.05 คัน และ 0.14 คัน ตามลำดับ

4.3) การจัดแถวคอยที่ 3

4.3.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท B, C และ D ลดลง 1.31 นาที, 0.11 นาที และ 12.61 นาที ตามลำดับ แต่เวลารอคอยของรถประเภท A เพิ่มขึ้นเล็กน้อยคือ 0.27 นาที

4.3.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B และ D เพิ่มขึ้น 0.46 คัน, 0.12 คัน และ 0.15 คัน ตามลำดับ แต่จำนวนของรถประเภท C ลดลงเล็กน้อยคือ 0.02 คัน

4.4) วิเคราะห์ผลจากการปรับปรุงแถวคอย

การจัดแถวคอยทั้ง 3 แนวทาง สามารถทำให้เวลารอคอยของรถขนส่งส่วนใหญ่ลดลง แต่อย่างไรก็ตามจำนวนรถที่ออกจากระบบไม่แตกต่างจากสถานการณ์ก่อนปรับปรุงมากนักและมีรถขนส่งบางประเภทที่มีค่าลดลง เนื่องจากการจัดแถวคอยมีเป้าหมายเพื่อลดเวลารอคอย ดังนั้นเพื่อเพิ่มจำนวนรถที่ออกจากระบบ จึงนำไปสู่การปรับปรุงการจัดวางสินค้าเพื่อลดระยะเวลาลำเลียงภายในคลัง

4.5) การจัดการคลังสินค้าโดยการวิเคราะห์ FSN

4.5.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท A, B, C และ D ลดลง 4.08 นาที, 6.15 นาที, 2.83 นาที และ 8.00 นาที ตามลำดับ

4.5.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B, C และ D เพิ่มขึ้น 1.00 คัน, 0.80 คัน, 0.07 คัน และ 0.12 คัน ตามลำดับ

4.6) วิเคราะห์ผลจากการทำ FSN

การทำ FSN ส่งผลให้เวลารอคอยของรถทุกประเภทลดลง และสามารถเพิ่มจำนวนรถทุกประเภทที่ออกจากระบบตามเป้าหมาย แตกต่างจากการจัดแถวคอยที่สามารถเพิ่มจำนวนรถที่ออกจากระบบสำหรับรถบางประเภทเท่านั้น

4.7) การจัดแถวคอยที่ 1 - การวิเคราะห์ FSN

4.7.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท A, B, C และ D ลดลง 4.15 นาที, 4.73 นาที, 10.24 นาที และ 5.79 นาที ตามลำดับ

4.7.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B, C และ D เพิ่มขึ้น 0.92 คัน, 1.00 คัน, 0.32 คัน และ 0.27 คัน ตามลำดับ

4.8) การจัดแถวคอยที่ 2 - การวิเคราะห์ FSN

4.8.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท A, B, C และ D ลดลง 3.77 นาที, 5.00 นาที, 3.36 นาที และ 15.49 นาที ตามลำดับ

4.8.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B, C และ D เพิ่มขึ้น 0.95 คัน, 0.72 คัน, 0.13 คัน และ 0.30 คัน ตามลำดับ

4.9) การจัดแถวคอยที่ 3 - การวิเคราะห์ FSN

4.9.1) เวลาารอคอยเฉลี่ย: พบว่ารถประเภท A, B, C และ D ลดลง 3.14 นาที, 5.51 นาที, 2.76 นาที และ 14.36 นาที ตามลำดับ

4.9.2) จำนวนรถที่ออกจากระบบเฉลี่ย: รถประเภท A, B, C และ D เพิ่มขึ้น 1.00 คัน, 0.63 คัน, 0.07 คัน และ 0.26 คัน ตามลำดับ

4.10) วิเคราะห์ผลการปรับปรุงแถวคอยร่วมกับการทำ FSN

การจัดแถวคอยทั้ง 3 แบบ ร่วมกับการจัดการคลังสินค้าโดยวิเคราะห์ FSN ส่งผลให้เวลารอคอยของรถทุกประเภทลดลง อีกทั้งยังทำให้จำนวนรถทุกประเภทออกจากระบบเพิ่มขึ้น

เมื่อเปรียบเทียบการปรับปรุงแถวคอยร่วมกับการทำ FSN ทั้ง 3 วิธีการ พบว่าการจัดแถวคอยที่ 1 - การวิเคราะห์ FSN ให้ผลลัพธ์ในภาพรวมที่ดีที่สุด ทั้งเวลารอคอยที่ลดลง และจำนวนรถทั้งหมดที่ออกจากระบบที่เพิ่มขึ้น

ผลลัพธ์จากสถานการณ์ทั้งหมด พบว่าการปรับปรุงกระบวนการส่วนใหญ่ส่งผลให้เวลารอคอยของรถแต่ละประเภทลดลง จำนวนรถที่ออกจากระบบเพิ่มมากขึ้นจากระบบปัจจุบัน งานวิจัยนี้จะเลือกสถานการณ์ที่มีประสิทธิภาพมากที่สุด ได้แก่ การจัดแถวคอยที่ 1 - การวิเคราะห์ FSN ซึ่งเป็นสถานการณ์ที่ทำให้เวลารอคอยลดลงและสามารถทำให้จำนวนรถทั้งหมดที่ออกจากระบบโดยเฉลี่ยเพิ่มขึ้นมากที่สุด นั่นหมายถึงมีโอกาสการสร้างรายได้สูงสุด เป็นสถานการณ์ที่เหมาะสมสำหรับงานวิจัย ทั้งนี้จำนวนรถที่เพิ่มขึ้นจะนำมาวิเคราะห์เพื่อประมาณต้นทุน รายได้ และผลกำไรที่เกิดขึ้น ดังรายละเอียดในหัวข้อถัดไป

5) การประมาณต้นทุนและผลประโยชน์ที่ได้รับ

การประมาณต้นทุนและผลประโยชน์ที่ได้รับจะพิจารณาเฉพาะช่วงเวลาการณศึกษา 13.00–17.00 น. เท่านั้น จากสถานการณ์ที่ถูกเลือกพบว่าจำนวนรถที่ออกจากระบบเฉลี่ยในช่วงเวลาดังกล่าวต่อเดือนเพิ่มขึ้น แสดงดังตารางที่ 10

ตารางที่ 10 : จำนวนรถที่ออกจากระบบเฉลี่ยในหนึ่งเดือน

ประเภทรถ	สถานการณ์ปัจจุบัน	สถานการณ์ปรับปรุง
A	243	271
B	191	221
C	99	109
D	49	57
ผลรวม	582	658

ข้อมูลจากตารางที่ 10 พบว่าจำนวนรถขนส่งที่ออกจากระบบทั้งหมดเพิ่มขึ้น 76 คัน จำนวนรถที่ออกจากระบบจะถูกนำไปคำนวณต้นทุนและรายได้ต่อเดือน เฉพาะช่วงเวลาที่กำหนด 13.00–17.00 น. เพื่อแสดงให้เห็นว่าสถานการณ์ดังกล่าวเพิ่มโอกาสในการขายสินค้า จากปริมาณรถที่ออกจากระบบที่เพิ่มขึ้น

สินค้าภายในคลังถูกจำแนกเป็นสินค้าหลัก ได้แก่ เหล็กตัวซี เหล็กท่อน้ำ เหล็กท่อนเหลี่ยม และเหล็กกล่องแบน ในงานวิจัยประมาณราคาขายสำหรับ เหล็กตัวซี 24 บาท/กิโลกรัม เหล็กท่อน้ำ 25 บาท/กิโลกรัม เหล็กท่อนเหลี่ยม 25 บาท/กิโลกรัม เหล็กกล่องแบน 25 บาท/กิโลกรัม

การคำนวณรายได้จะพิจารณาน้ำหนักเฉลี่ยของสินค้าที่ลำเลียงขึ้นรถขนส่งร่วมกับราคาขายที่กำหนด สามารถหารายได้รวมต่อคันสำหรับรถแต่ละประเภท แสดงดังตารางที่ 11

ตารางที่ 11 : รายได้รวมต่อรถขนส่ง 1 คัน

ประเภทรถ	ชนิดสินค้า	ปริมาณสินค้าเฉลี่ย (กิโลกรัม/คัน)	รายได้รวม/คัน (บาท)
A	เหล็กตัวซี	2114.83	154505.16
	เหล็กท่อน้ำ	1230.77	
	เหล็กท่อนเหลี่ยม	1385.44	
	เหล็กกล่องแบน	1533.77	
B	เหล็กตัวซี	3203.63	228675.47
	เหล็กท่อน้ำ	1825.16	
	เหล็กท่อนเหลี่ยม	1918.62	
	เหล็กกล่องแบน	2327.75	
C	เหล็กตัวซี	1398.45	109898.78
	เหล็กท่อน้ำ	983.33	
	เหล็กท่อนเหลี่ยม	971.41	
	เหล็กกล่องแบน	1098.69	
D	เหล็กตัวซี	2873.55	222205.11
	เหล็กท่อน้ำ	1555.63	
	เหล็กท่อนเหลี่ยม	1836.64	
	เหล็กกล่องแบน	2737.33	

การวิเคราะห์ต้นทุนจะพิจารณาข้อมูลปริมาณเหล็กทุกประเภทภายในคลัง พบว่าต้นทุนทั้งหมด คือ 57,862,590.46 บาท และเมื่อพิจารณาจำนวนรถที่ออกจากระบบ พบว่ารายได้เฉลี่ยต่อเดือนและผลกำไรเฉลี่ยต่อเดือน ก่อนและหลังปรับปรุงกระบวนการแสดงดังตารางที่ 12 ซึ่งพบว่าหลังการปรับปรุงกระบวนการนอกจากช่วยลดระยะเวลาการคอยเฉลี่ยในแถวคอยและเพิ่มจำนวนรถโดยเฉลี่ยที่ออกจากระบบแล้วยังเพิ่มกำไรให้สถานประกอบการ 14,063,037.11 บาท/เดือน

ต้นทุน รายได้ และผลกำไรเป็นการประเมินจากผลลัพธ์ที่เกิดขึ้นจากแบบจำลองทางคอมพิวเตอร์ นั่นคือปริมาณเหล็กที่อยู่ในคลังจากการเก็บข้อมูลและปริมาณเหล็กที่ขายได้ซึ่งประเมินจากจำนวนรถที่ออกจากระบบ ทั้งนี้หากจะนำแนวคิดการปรับปรุงงานจากแบบจำลองทางคอมพิวเตอร์ไปดำเนินการปรับปรุงระบบงานจริงจะต้องคำนึงถึงต้นทุนที่เกี่ยวข้องกับการเปลี่ยนแปลงการทำงานอื่น ๆ ร่วมด้วย เช่น ต้นทุนในการจัดระเบียบสินค้าใหม่

ภายในคลัง ต้นทุนค่าแรงงานสำหรับการเปลี่ยนแปลงระบบงาน ต้นทุนค่าเสียโอกาสในช่วงเปลี่ยนผ่านกระบวนการทำงาน เป็นต้น

ตารางที่ 12 : ต้นทุน รายได้ และผลกำไรโดยเฉลี่ยต่อเดือน

ต้นทุน/รายได้/กำไร	สถานการณ์ปัจจุบัน	สถานการณ์ปรับปรุง
ต้นทุนเฉลี่ย (บาท)	57,862,590.46	57,862,590.46
รายได้เฉลี่ย (บาท)	102,989,797.22	117,052,834.33
กำไรเฉลี่ย (บาท)	45,127,206.76	59,190,243.87

6) สรุปผล

งานวิจัยนี้ศึกษากระบวนการขึ้นสินค้าของบริษัทกรณีศึกษา โดยเริ่มจากการศึกษากระบวนการปัจจุบัน พบว่ามีปัญหาความล่าช้าเกิดขึ้นในกระบวนการในช่วงเวลา 13.00–17.00 น. แบ่งออกเป็น 2 กรณี ได้แก่ กรณีที่ 1 รถขนส่งมีเวลารอคอยนานก่อนจะเข้าไปรับสินค้า กรณีที่ 2 การลำเลียงสินค้าภายในคลังเพื่อนำขึ้นรถขนส่งใช้เวลานาน ปัญหาดังกล่าวเกิดจากบริษัทกรณีศึกษายังไม่มีกระบวนการจัดการแฉกคอคยที่เหมาะสม และยังไม่มีรูปแบบสำหรับการจัดวางสินค้าภายในคลัง (5 เฟส) เมื่อทราบถึงขั้นตอนการทำงานปัจจุบันแล้วจึงเก็บรวบรวมข้อมูลที่เกี่ยวข้องกับการดำเนินงานภายในคลัง และนำไปสร้างแบบจำลองสถานการณ์ทางคอมพิวเตอร์เพื่อใช้เป็นตัวแทนของกระบวนการทำงานปัจจุบัน

หลังจากตรวจสอบความถูกต้องและความเสมือนจริงของแบบจำลองแล้วจึงกำหนดแนวทางการปรับปรุงระบบงานผ่านแบบจำลองทางคอมพิวเตอร์ โดยแบ่งออกเป็นการจัดแฉกคอคย 3 แนวทาง และการวิเคราะห์พื้นที่จัดวางสินค้าภายในเฟสโดยใช้การวิเคราะห์ FSN ทำให้ได้แนวทางการปรับปรุงทั้งหมด 7 สถานการณ์ แต่ละสถานการณ์ประมวผล 35 รอบทำซ้ำ ในช่วงเวลา 13:00 น.-17:00 น. โดยใช้เวลารอคอย และจำนวนรถที่ออกจากกระบบ เป็นดัชนีชี้วัดประสิทธิภาพของกระบวนการ

ผลการจำลองสถานการณ์ปรับปรุงของ 7 แนวทางเปรียบเทียบกับสถานการณ์ก่อนปรับปรุง พบว่าการจัดการแฉกคอคย 3 แนวทาง สามารถทำให้เวลารอคอยของรถขนส่งส่วนใหญ่ลดลง ส่วนการจัดสรรพื้นที่จากการทำ FSN ทำให้เวลารอคอยของรถขนส่งลดลงและยังทำให้รถขนส่งทุกประเภทที่ออกจากกระบบเพิ่มมากขึ้น อย่างไรก็ตามเมื่อนำการจัดการแฉกคอคยมาใช้ร่วมกับการวิเคราะห์ FSN ยังพบว่าเวลารอคอยของรถขนส่งทุกประเภทลดลง อีกทั้งจำนวนรถขนส่งของรถทุกประเภทที่ออกจากกระบบเพิ่มขึ้น

ในงานวิจัยเลือกการจัดแฉกคอคยที่ 1 - การวิเคราะห์ FSN เป็นสถานการณ์ที่เหมาะสมสำหรับงานวิจัย เนื่องจากสามารถลดเวลารอคอยเฉลี่ยของรถทุกประเภท และทำให้จำนวนรถที่ออกจากกระบบโดยเฉลี่ยเพิ่มขึ้นในภาพรวมมากที่สุดเมื่อเปรียบเทียบกับสถานการณ์อื่น ๆ สถานการณ์ดังกล่าวเป็นสถานการณ์ที่สร้างโอกาสในการขายสินค้าและเพิ่มกำไรให้แก่บริษัทกรณีศึกษามากที่สุด

จากผลการวิจัยแสดงให้เห็นว่าสถานการณ์ที่เหมาะสมและสามารถลดความล่าช้าของกระบวนการขึ้นสินค้าได้ดีที่สุดเกิดจากการนำการจัดระเบียบแฉกคอคยและการจัดวางสินค้าจากการวิเคราะห์ FSN มาใช้ร่วมกัน เนื่องจากปัญหาในงานวิจัยเกิดในการดำเนินงานที่ต่อเนื่องกัน คือ เวลารอคอยของรถขนส่งในแฉกคอคยนานก่อนที่รถขนส่งจะเข้าไปรับสินค้าภายในคลังซึ่งมีเวลาการลำเลียงสินค้ายาวนาน การผสมผสาน 2 เทคนิคเข้าด้วยกันจึงเป็นสถานการณ์ที่เกิดประสิทธิภาพสูงที่สุดในงานวิจัย

7) ข้อเสนอแนะ

จากการวิจัยมีข้อสังเกตซึ่งเป็นสิ่งที่น่าสนใจทำการศึกษาและวิจัยต่อไปในอนาคต ดังนี้

หากพิจารณาการจัดการแฉกคอคยที่ 1 - การวิเคราะห์ FSN ซึ่งเป็นวิธีการที่ถูกเลือก จะเห็นว่าเวลารอคอยของรถประเภท B และ D มากกว่าการจัดการแฉกคอคยร่วมกับการทำ FSN อีก 2 วิธี สามารถวิเคราะห์ได้ว่าแนวทางปรับปรุงกระบวนการดังกล่าวเป็นการให้ความสำคัญกับรถประเภท C ดังนั้น อาจทำให้รถประเภทอื่นๆ ซึ่งมีช่วงห่างการมาถึงมากกว่าหรือใกล้เคียงมีเวลารอคอยที่เพิ่มขึ้นเล็กน้อยเนื่องจากรถประเภท C จะมีสิทธิ์เข้าคลังก่อนเสมอเมื่อมาถึงระบบงาน ดังนั้นเมื่อเปรียบเทียบกับแนวทางปรับปรุงอื่นๆ เวลารอคอยจะมากกว่าเล็กน้อย อย่างไรก็ตามจากผลการจำลองสถานการณ์พบว่าเวลารอคอยของรถประเภท B และ D ที่เพิ่มขึ้นเล็กน้อยจากสถานการณ์ปรับปรุงอื่นๆ ไม่ส่งผลกระทบต่อระบบ เนื่องจากยังเป็นแนวทางการปรับปรุงที่ทำให้ผลรวมจำนวนรถที่ออกจากกระบบเฉลี่ยมากที่สุด จึงถูกเลือกให้เป็นสถานการณ์ที่เหมาะสม สามารถสร้างรายได้และผลกำไรได้สูงที่สุด ทั้งนี้การหาแนวทางพัฒนาระบบเพื่อให้เวลารอคอยของรถทุกประเภทต่ำที่สุดซึ่งจะส่งผลให้จำนวนรถที่ออกจากกระบบเพิ่มมากขึ้น เป็นสิ่งที่น่าสนใจศึกษาทำให้ประสิทธิภาพของระบบงานเพิ่มขึ้น

และจากการสังเกตแบบจำลองทางคอมพิวเตอร์พบว่าหลังการปรับปรุงกระบวนการขึ้นสินค้า เฟสว่างบ่อขึ้น แสดงให้เห็น

ถึงค่าอรรถประโยชน์มีแนวโน้มลดลง เนื่องจากเวลารอคอยของรถขนส่งลดลงและเวลาลำเลียงสินค้าเร็วขึ้น ทำให้เกิดช่วงเวลาที่ไม่มีรถเข้ามาในเฟสบ่อยขึ้น เฟสที่ว่างมากขึ้นแสดงให้เห็นว่าคลังสินค้ามีศักยภาพในการรองรับรถขนส่งที่มากขึ้นได้ หากสามารถบริหารจัดการให้อัตราการเข้ามาของรถขนส่งเพิ่มขึ้น จะช่วยให้โรงงานกรณีศึกษาสามารถเพิ่มอรรถประโยชน์ของเฟสขนส่งสินค้าได้เต็มที่และเพิ่มโอกาสการสร้างกำไรให้กับบริษัทได้มากขึ้น

นอกจากนี้เมื่อพิจารณาพื้นที่การจัดเก็บสินค้า พบว่าปริมาณสินค้าประเภท FSN ที่อยู่ในแต่ละเฟสมีปริมาณไม่แน่นอน ทำให้สินค้าไม่ครบถ้วนตามใบสั่งซื้อ จึงต้องมีค่าเผื่อเวลาการลำเลียงสินค้าจากเฟสอื่น จึงเป็นเหตุผลอีกประการหนึ่งที่ทำให้การขึ้นสินค้าเกิดความล่าช้า หากปรับสัดส่วนการจัดเก็บให้เหมาะสมและสามารถระบุตำแหน่งจัดวางได้ละเอียดมากขึ้น จะสามารถลดเวลาการลำเลียงสินค้าลงและทำให้รูปแบบการลำเลียงมีประสิทธิภาพมากขึ้น

การพัฒนากระบวนการลดเวลารอคอย การวิเคราะห์อรรถประโยชน์ และปรับสัดส่วนการจัดเก็บให้เหมาะสมต้องมีการเก็บข้อมูลและวิเคราะห์ข้อมูลเพิ่มเติม ซึ่งเป็นสิ่งที่น่าสนใจสำหรับการวิจัยต่อไปในอนาคต

กิตติกรรมประกาศ

ผู้วิจัยขอขอบคุณ บริษัท ทีเอ็มที สตีล จำกัด (มหาชน) ที่ให้ความอนุเคราะห์ข้อมูลสำหรับการสร้างแบบจำลองทางคอมพิวเตอร์ รวมถึงข้อเสนอแนะที่เป็นประโยชน์ ทำให้การพัฒนางานวิจัยนี้เกิดประโยชน์สูงสุด

REFERENCES

- [1] Kasikorn Research Center. "Thai steel prices projected to fall by 6% YoY in 2024 due to intense competition from imported steel from China," *Current Issue*, no. 3475, Mar. 2024. [Online] Available: <https://www.kasikornresearch.com/EN/analysis/k-econ/business/Pages/Steel-CIS3475-23-03-2024.aspx>
- [2] J. Banks, J. S. Carson, B. L. Nelson, and D. M. Nicol, *Discrete-Event System Simulation*, 4th ed. Upper Saddle River, NJ, USA: Prentice Hall. 2005.
- [3] R. Pisuchpen, *The Handbook of Simulation Modeling with ARENA*, revised ed. Bangkok, Thailand: SE-EDUCATION (in Thai), 2010.
- [4] J. Pichitlamken, *Fundamentals of Random Simulation for Real-World Applications*. Bangkok, Thailand: Kasetsart University Press (in Thai), 2015.
- [5] D. Devarajan and M. S. Jayamohan, "Stock control in a Chemical Firm: Combined FSN and XYZ Analysis," *Procedia Technol.*, vol. 24, pp. 562–567, 2016.
- [6] D. C. Montgomery and G. C. Runger, "Hypothesis testing for two populations," in *Engineering Statistics*, P. Sudasna-Na-Ayudhya and P. Luangpaiboon, Eds., 3rd ed. Bangkok, Thailand: Top Publishing (in Thai), 2011, ch. 8, sec. 2, pp. 251–266.
- [7] X. Zhu, R. Zhang, F. Chu, Z. He, and J. Li, "A FlexSIM-based optimization for the operation process of cold-chain logistics distribution centre," *J. Appl. Res. Technol.*, vol. 12, no. 2, pp. 270–278, 2014.
- [8] M. A. Millstein, L. Yang, and H. Li, "Optimizing ABC inventory grouping decisions," *Int. J. Prod. Econ.*, vol. 148, pp. 71–80, 2014.
- [9] N. Panyapanich and S. Sirisoponsilp, "The application of a simulation model to analyze the efficiency of order picking process," (in Thai), *Chulalongkorn Bus. Rev.*, vol. 37, no. 144, pp. 24–50, 2015.
- [10] N. Wongsamajan, "Finish goods storage layout improvement to reduce the total inventory movement: a case of foundry and machining manufacturing," (in Thai), M.S. independent study, Dept. Ind. Eng., Thammasat Univ., Pathum Thani, Thailand, 2017.
- [11] C. Bunterngchit, "Simulation-based application in warehouse layout design for reducing material handling time," (in Thai), *Kasem Bundit Eng. J.*, vol. 8, no. 3, pp. 1–14, 2018.
- [12] T. Chanyeam and N. Akarapichet, "Reduction the delivery time spare parts for improvement the spare parts lay out: Case study of AY Company Limited," (in Thai), in *Proc. 6th Nat. Conf. Nakhonratchasima Coll.*, Nakhon Ratchasima, Thailand, Mar. 2019, pp. 232–250.
- [13] C. Anurattananon, P. Klomjit, T. Chuenyindee, and P. Songsuktawan, "Efficiency improvement of drinking warehouse case study of sample drinking company," (in Thai), *Thai Ind. Eng. Netw. J.*, vol. 5, no. 1, pp. 49–58, 2019.
- [14] P. Petchan and A. Karoonsoontawong, "Development of simulation for queuing system at Saphan-Taksin BTS station by using PTV Viswalk," (in Thai), *Kasem Bundit Eng. J.*, vol. 10, no. 3, pp. 25–45, 2020.

- [15] C. Kraiwong and S. Setamanit, "Loading process improvement of wood substitute industry using value stream mapping and simulation," (in Thai), *Res. J. Rajamangala Univ. Technol. Thanyaburi*, vol. 19, no. 2, pp. 60–72, 2020.
- [16] O. Ganbold, K. Kundu, H. Li, and W. Zhang, "A Simulation-based optimization method for warehouse worker assignment," *Algorithms*, vol. 13, no. 12, 2020, Art. no. 326.
- [17] P. Jirapatsil, "Improvement of fresh food location and outbound process for e-commerce business warehouse," (in Thai), M.S. thesis, Dept. Ind. Eng., Chulalongkorn Univ., Bangkok, Thailand, 2021.
- [18] A. Chaimanee and S. Kaewcharoensithong, "The simulation to reduce rice in the process of paddy processing," (in Thai), *UBU Eng. J.*, vol. 14, no. 4, pp. 24–37, 2021.
- [19] P. Preedawat, W. Panprung, and W. Atthirawong, "Improving efficiency of warehouse management using a simulation technique: A case study of industrial gas distribution company," (in Thai), *J. Appl. Stat. Inf. Technol.*, vol. 7, no. 2, pp. 26–47, 2022.
- [20] K. Phimkan, "The management of hardware warehouse maintenance system of Ordnance Company Anti-aircraft Artillery Division using ABC-FSN analysis," (in Thai), *CRMA J.*, vol. 20, pp. 62–74, 2022.
- [21] P. Likitkarn, "Improving the inventory storage process with ABC analysis for increase storage management efficiency: The case of XYZ Co., Ltd.," (in Thai), *J. Logist. Supply Chain Coll.*, vol. 9, no. 1, pp. 5–19, 2023.
- [22] C. Nigischer, F. Reiterer, S. Bougain, and M. Grafinger, "Finding the proper level of detail to achieve sufficient model fidelity using FlexSim: An industrial use case," *Procedia CIRP*, vol. 119, pp. 1240–1245, 2023.
- [23] K. Khotyota, W. Ngamsaard, and P. Nakseedee, "Reducing picking errors and picking time in warehouse: Case study of Eurosia Foods Trading & Agencies Co., Ltd.," (in Thai), *J. Sci. Technol., Southeast Bangkok Univ.*, vol. 3, no. 2, pp. 14–31, 2023.
- [24] S. Tuammee, S. Cheevapruk, P. Rontlaong, P. Junsanad, and T. Pikul, "Efficiency improvement of location assignment of products in a warehouse: An empirical study," (in Thai), in *Proc. 14th Hatyai Nat. and Int. Conf.*, Songkhla, Thailand, May 2023, pp. 643–658.
- [25] T. Sringean, P. Nakseedee, and W. Ngamsaard, "Increase inventory management performance: A case study of a government agenc AAA," (in Thai), *J. Sci. Technol., Southeast Bangkok Univ.*, vol. 3, no. 1, pp. 13–28, 2023.
- [26] K. Chatsatayu, C. Budsupho, S. Doemsriphum, S. Sajjanan, S. Dumsrikong, and K. Piriyaipun, "The improvement of the system and storage location to reduce time for finding and picking products in packaging warehouses: A case study of ABC Co., Ltd.," (in Thai), *Sripatum Chonburi Interdiscip. J. (Online)*, vol. 9, no. 2, pp. 30–47, 2023.
- [27] D. P. Ardiansyah, F. A. O. Reynaldi, and D. L. Widaningrum, "Improving warehouse layout effectiveness and process picking efficiency with the discrete event system simulation approach," *Procedia Comput. Sci.*, vol. 234, pp. 1753–1760, 2024.
- [28] P. Lumsakul, P. Saengkhiao, and P. Kaweejitbundit, "Improvement of inbound logistics process in coconut manufacturing using FlexSim simulation: A case study," *J. Eng. Digit. Technol. (JEDT)*, vol. 12, no. 1, pp. 12–22, 2024.
- [29] R. G. Sargent, D. M. Goldsman, and T. Yaacoub, "A tutorial on the operational validation of simulation models," in *Proc. 2016 Winter Simulation Conf. (WSC)*, Washington, DC, USA, Dec. 2016, pp. 163–177.

Optimal Allocation and Deployment of Roadside Units in Cloud-Based Internet of Vehicles Framework

Nay Myo Sandar^{1*} Surekha Lanka² Thinzar Aung Win³ Shuvra Tripura⁴

^{1*,2,3,4}*Faculty of Business and Technology, Stamford International University, Bangkok, Thailand*

*Corresponding Author. E-mail address: naymyo.sandar@stamford.edu

Received: 13 February 2025; Revised: 4 June 2025; Accepted: 8 June 2025

Published online: 26 December 2025

Abstract

The research focuses on internet of vehicles (IoV) where vehicles are equipped with cameras and sensors to monitor traffic jams, accidents, and locations to ensure safety and comfort for drivers. To process sensor data effectively, cloud computing is used because of vast storage and processing capabilities. However, transferring data from sensors to cloud can be challenging due to bandwidth and memory constraints. Therefore, a cloud-based internet of vehicles framework is proposed incorporating Roadside Units (RSUs). RSUs can buffer video streams from vehicles and send them to cloud services. With RSUs in the framework, total latency for transferring video streams to cloud services can be significantly enhanced. In this research, dynamic programming approaches are applied to determine how many RSUs are needed at the lowest cost and greedy algorithm is implemented to prove the optimal solution from dynamic programming. Furthermore, K-means clustering algorithm is applied to find the best locations for RSUs. According to numerical results, the proposed methods can determine the optimal number of 6 RSUs with the minimum cost and allocation of RSUs to serve video streams across regions.

Keywords: Cloud-based internet of vehicles, Dynamic programming approach, Greedy algorithm, K-means clustering algorithm, Roadside units (RSUs)



I. INTRODUCTION

Today, the number of vehicles used in many countries has exceeded 1 billion [1], largely driven by economic and population growth, and is expected to reach 2 billion by 2035 [2]. Due to the increasing number of vehicles, it can raise traffic congestion and accidents on the roads. Until recently, the internet of vehicles (IoV) is emerged to solve the challenging issues in current transportation systems [3]. In IoV system, the intelligent sensor devices are installed in vehicles and traffic lights which can be able to absorb data from the environment and exchange information between vehicles to vehicles and vehicles to roads [4]. The use of IoV system can provide vehicle transportation safety and improve users' convenience [5]. However, data readings from sensors installed in vehicles and traffic lights need to be processed to become useful and meaning information. As a result, cloud computing has been emerged to provide scalable storage capacity and processing power for achieving value-added services such as data analytics, backup, database management system, respectively [6]. Therefore, data streamed from sensors can be transferred to the cloud for storing and processing in order to utilize value-added services. Then, data processed or stored by the cloud services can be remotely monitored and managed by the users through their Internet connected devices such as surveillance monitors, personal computers, smart phones, tablets, and set-top boxes.

However, although cloud computing provides a large pool of resources to use cloud services, serious network latency could be encountered when sensors with limited bandwidth transfer data to the cloud services over long distance [7]. As a result, the small amount of buffer memory allocated in sensors will be overflow and data can be lost. To address such problem, Cloud-based Internet of Vehicles framework in which stream aggregation (SA) approach is applied. The SA approach provides Roadside Units (RSUs) which are buffering

sensors with high speed network connectivity between sensors and cloud services [8]. An RSU can be allocated close to vehicles on the roads and collect the data from sensors and move to the cloud services with better bandwidth. In such case, RSUs are manufactured by some companies Savari, Fluidmesh Networks, Beijing Juli Science & Technology Co., Ltd., etc. [9].

Since the framework is done by RSUs, it is challenging to determine how many RSUs and where should be located. For the proposed issues, we applied dynamic programming optimization approach to minimize the total cost for allocating the optimal number of RSUs. Dynamic programming is problem solving technique which breaks down problem into simpler subproblems, solves each by one, stores the results, and finally combines the results [10]. After that, greedy algorithm is applied to prove the optimal number of RSUs found by dynamic programming. A greedy algorithm is problem solving approach which makes the best choice in every stage and finally lead to overall optimal solution [11]. Moreover, K-means clustering algorithm is utilized to find the optimal location for installing RSUs in areas of Internet of vehicles in the proposed framework. This algorithm groups data into clusters by assigning to the nearest centroid [12]. Then, we evaluate the proposed optimization approach and algorithms by performing numerical studies.

This research paper presents literature review in Section 2, research methods in Section 3, results in Section 4, and conclusion and discussion in Section 5.

II. LITERATURE REVIEW

Recently, earlier researchers have studied the combination of cloud computing and internet of vehicles. For example, a novel multilayered vehicular data cloud platform was proposed in [13] by using cloud computing and internet of vehicles to store and process transportation related information such as

traffic control and management, car location tracking and monitoring, road condition, car warranty, and maintenance information for drivers. In [14], edge-cloud computing is deployed for autonomous vehicles to provide a large amount of computation resources. Cloud-assisted Internet of Things Intelligent Transportation System (CIoT-ITS) was proposed in [15] to handle traffic management's challenges.

Based on the papers, cloud computing has integrated with the internet of vehicles since it can provide the benefits of virtual servers that allow remote storage and computational capacities. Data stored in clouds can be retrieved from anywhere and can be shared among the vehicles. However, one critical issue of network scalability problem needs to be considered when many vehicles transfer data to cloud services over long range communication.

Dealing with these problems, the concept of roadside units (RSUs) is introduced. RSUs can perform as an aggregator which collects data from vehicles and transfers to the cloud for further storage and processing. In [16], aggregators are applied in a smart (electrical) grid to achieve scalable connectivity which can transfer data about metered power usage from many smart meters to the same target. In [17], a smart QoE-Aware Adaptive Video Bitrate Aggregation scheme is proposed for HTTP live streaming on smart edge computing to improve efficiency of live video streaming. In [18], the researchers presented roadside units (RSUs) which collect traffic data from vehicles and transmit them to the traffic control center. The article in [19] proposed intelligent transportation system with roadside unit (RSU) using the concept of internet of things. Based on the concepts established in previous research, this paper applies cloud computing technology and roadside units to enhance connectivity for Internet-connected vehicles.

Although previous studies have proposed cloud computing technology and the use of roadside units for the Internet of Vehicles, they have not addressed the unified challenge of simultaneously optimizing the number of roadside units and determining optimal deployment locations for roadside units. Different from the previous works, this paper utilizes dynamic programming approach for roadside unit (RSU) allocation, and the K-means clustering algorithm for RSU deployment to minimize the overall cost while meeting the demand of video streams within the proposed framework.

III. RESEARCH METHODS

In this section, we initially proposed the cloud-based internet of vehicles framework for traffic monitoring system as shown in Figure 1 where the bottom layer is internet of vehicles (IoV) that is a self-organized network by connecting vehicles which are equipped with wireless sensors in order to collect road data such as traffic jam, accident detection, post-accident investigation, and intersection collision, respectively. The collection of large amounts of road data needs to be processed to provide useful information for a wide variety of services (e.g., transport safety, traffic management, convenient driving, etc.). However, vehicle sensors are normally constrained by resources of storage and computation so that it becomes difficult to achieve those services. As a result, cloud computing is a very promising technology which can offer an excessive amount of storage resources and computation resources. By integrating cloud computing to the internet of vehicles, road data could be collected by vehicles and transmitted to the cloud for processing and in return, information could be broadcasted to vehicles for better traffic control and safety on the road.

In the framework, vehicles transmit road data to the traffic lights beside the road. However, traffic lights have limited wireless bandwidth so that undesirable network

latency can be experienced for uploading data to the cloud. For this issue, Roadside Units (RSUs) are further installed between traffic lights and cloud infrastructure which can perform as stream aggregation approach which collects the data reported by traffic lights and send to the cloud for further analyzing and processing. After the cloud processes data, it can provide traffic information which can help drivers to choose the optimal routes by avoiding potential accidents, dangerous situations, and traffic congestion along the road.

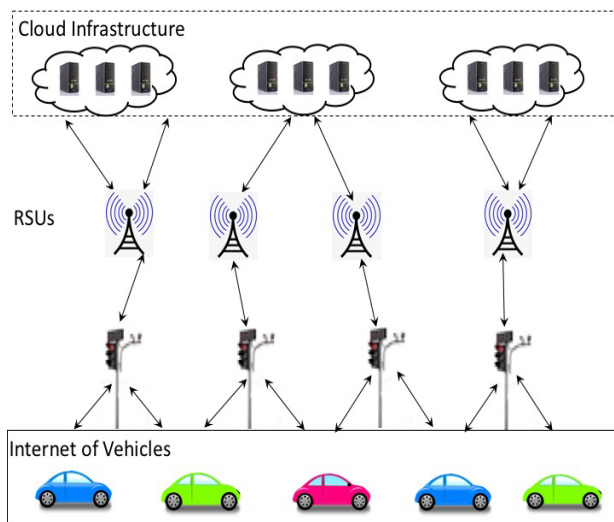


Figure 1: Cloud-based Internet of vehicles framework

Since RSUs are applied in the framework, the installation of RSUs is very expensive and it is difficult to determine the number of required RSUs to cover several vehicles on the road. Additionally, deploying RSUs along the road to ensure coverage of vehicles is also challenging. To solve these issues, a dynamic programming optimization approach is applied to obtain the optimal number of RSUs in the framework. Then, we implemented greedy algorithms to verify the optimal number of RSUs obtained by dynamic programming. Furthermore, K-means clustering algorithm is also applied to determine in which locations for deploying RSUs on the roadside. Then, we perform numerical studies to evaluate the optimal solutions.

A. Dynamic Programming Approach to Roadside Units Allocation

First, a dynamic programming approach is applied for determining the optimal number of Roadside Units (RSUs) with the minimum cost needed to handle video streams from traffic lights on the road. Each RSU has a specific amount of deployment cost and has maximum capacity to handle up to a certain number of video streams. Let streams $[i]$ represent the number of video streams generated by traffic lights i . Let $dp[i]$ is dynamic programming array which represents the minimum number of RSUs needed to handle video streams from traffic lights i .

B. Greedy Algorithm Approach to Prove Roadside Units Allocation

Greedy algorithm is also applied to prove the optimal number of Roadside Units (RSUs) obtained by dynamic programming approach.

Table 1: Greedy algorithm for optimal allocation of RSU

Input:
Video Streams: List of stream sizes
Max Capacity: Maximum capacity per RSU
Output:
Number of RSUs: Total RSU used
Step 1: Initialize variables
RSUs = []
currentRSUCapacity = 0
Step 2: Sort video streams in descending order
sortedStreams = sort (videoStreams, descending)
Step 3: Iterate through video streams
if currentRSUCapacity + stream <= maxCapacity:
currentRSUCapacity += stream
else: RSUs.append(currentRSUCapacity)
currentRSUCapacity = stream
Step 4: Add the last RSU capacity if not empty
if currentRSUCapacity > 0:
RSUs.append(currentRSUCapacity)
Return the total number of RSUs allocated
Return length (RSUs)

As shown in Table 1, The greedy algorithm involves four main steps. First, the algorithm checks for each stream if it can accommodate into the existing RSU without exceeding the maximum capacity. Second, if the stream can fit into existing RSU without exceeding its capacity, it is added to that RSU. This is greedy because it utilizes the available resources of existing RSU before creating new ones. Third, if the stream doesn't fit into any existing RSU, the algorithm creates a new RSU for that stream. This algorithm iterates through each vehicle's streams from all traffic lights. After iterating through all video streams from all traffic lights, it calculates the total number of RSUs needed and the total cost.

C. K-means Clustering Algorithm for Roadside Units Deployment

Next, we implement K-means clustering algorithm as shown in Table 2 to find the optimal locations for installing RSUs along the roadway to meet video stream demands. Each RSU needs to cover specific number of video streams from traffic lights with minimal distance to improve connectivity.

Table 2: K-means clustering algorithm for optimal placement of RSU

Input: Traffic Lights = [{index: 0, demand: 0}, {index: 1, demand: 3}, ...]
Step 1: Initialize RSUs at random traffic light positions Initialize K RSU centers at random traffic light indices
Repeat until convergence:
Step 2: Assign traffic lights to closest RSU
For each traffic light:
Find nearest RSU with remaining capacity
Assign traffic light if RSU capacity allows
Update RSU's remaining capacity
Step 3: Update RSU centers to centroid of assigned traffic lights
For each RSU with assigned traffic lights:
Set RSU center to average position of its traffic lights
End Repeat
Output: Final RSU positions and assigned traffic lights

The input Traffic Lights array contains each traffic light's index and demand of video streams. We start placing K RSUs at random traffic light positions along the road. These initial placements serve as RSU locations which will be adjusted in each iteration to cover traffic light demands. For each traffic light, the algorithm calculates the nearest RSU based on Euclidean distance. If an RSU has enough remaining capacity to serve the traffic light's demand, the traffic light is assigned to it. After assignment, the RSU's remaining capacity is updated for the new demand. If an RSU cannot accommodate to a traffic light due to capacity constraints, the algorithm tries to assign to the next closest RSU with available capacity. Once all traffic lights are assigned, each RSU is updated to the centroid of all traffic lights associated with that RSU, ensuring that each RSU center is optimally located to handle the data demands from traffic lights while minimizing travel distance between each RSU and traffic lights.

IV. RESULTS

First, we calculate the minimum number of roadside units (RSUs) required to handle video streams from 10 traffic lights using the dynamic programming approach. It is assumed that there are 10 video streams from 10 traffic lights streams[i] = [3,2,4,1,5,2,3,4,1,2]. Each RSU costs \$100 and it can handle up to 5 video streams.

For each traffic light i, the minimum number of RSUs required to process streams from traffic lights i can be calculated using the formula:

$$dp_i = \min(dp[j] + 1), \text{ where } \sum_{k=j}^i S_k \leq C \quad (1)$$

- $dp[i]$ represents the minimum number of RSUs needed to handle the video streams from traffic light i.
- $dp[j] + 1$ accounts for the new RSU to cover video streams from previous traffic light j to traffic light i.

- $\sum_{k=j}^i S_k \leq C$ ensures that the number of streams from traffic light j to i does not exceed the capacity of a single RSU. If this condition is met, then a single RSU can handle the streams from traffic light j to traffic light i.

Then, cumulative streams are calculated to get the number of streams between any two traffic lights.

$$\text{streams}[i] = [3,2,4,1,5,2,3,4,1,2] \quad (2)$$

$$\text{Cumulative_streams} = [3,5,9,10,15,17,20,24,25,27]$$

Based on the numerical results in Table 3, 6 RSUs are needed to process the video streams from 10 traffic lights. After we obtain the required number of RSUs, we calculate the minimum total cost using the formula:

$$\text{Total Cost} = \text{RSUs} \times \text{Cost per RSU} \quad (3)$$

$$\text{Total Cost} = 6 \times \$100 = \$600$$

The minimum total cost for installing the optimal number of RSU required to process the video streams from traffic lights is \$600. Then, we prove this optimal solution using greedy algorithm. Greedy algorithm is implemented in Java program to allocate videos streams from traffic lights to RSUs efficiently. The goal is to minimize the number of RSUs needed while ensuring each RSU does not exceed its capacity. Figure 2 shows the generated output of greedy algorithm.

Traffic Light Index	Video Streams	RSU Allocation Status
1	3	New RSU 1 (Total: 3)
2	2	Added to RSU 1 (Total: 5)
3	4	New RSU 2 (Total: 4)
4	1	Added to RSU 2 (Total: 5)
5	5	New RSU 3 (Total: 5)
6	2	New RSU 4 (Total: 2)
7	3	Added to RSU 4 (Total: 5)
8	4	New RSU 5 (Total: 4)
9	1	Added to RSU 5 (Total: 5)
10	2	New RSU 6 (Total: 2)

RSU Management Summary
Total RSUs needed: 6
Total cost: \$600

Figure 2: Numerical results of greedy algorithm

As shown in Figure 2, traffic light indexes 1, 3, 5, 6, 8, and 10 require new RSUs because the streams cannot be accommodated into the existing RSUs. Traffic light indexes 2, 4, 7, and 9 can add streams to existing RSUs to utilize the available resources. Finally, the program displays the same numerical results that 6 RSUs are required with the overall total cost of deployment of \$600.

To determine the areas where 6 RSUs should be installed based on the vehicle stream data, K-means clustering approach is applied. Based on the K-means clustering approach, we distribute the traffic light indices into clusters and each cluster represents an RSU location along the road. The clustering result for 6 RSUs is presented in Table 4 (in the next page).

Table 3: Numerical results for optimal number of RSUs

Traffic Light Index i	Video Streams	Cumulative Streams	Formula	Calculation	RSUs Needed
1	3	3	$dp[1]=dp[0]+1$	$dp[1]=0+1=1$	1 (Fits in RSU 1)
2	2	5	$dp[2]=dp[1]+0$	$dp[2]=1+0=1$	1 (Fits in RSU 1)
3	4	9	$dp[3]=dp[2]+1$	$dp[3]=1+1=2$	2 (New RSU needed)
4	1	10	$dp[4]=dp[3]+0$	$dp[4]=2+0=2$	2 (Fits in RSU 2)
5	5	15	$dp[5]=dp[4]+1$	$dp[5]=2+1=3$	3 (New RSU needed)
6	2	17	$dp[6]=dp[5]+1$	$dp[6]=3+1=4$	4 (New RSU needed)
7	3	20	$dp[7]=dp[6]+0$	$dp[7]=4+0=4$	4 (Fits in RSU 4)
8	4	24	$dp[8]=dp[7]+1$	$dp[8]=4+1=5$	5 (New RSU needed)
9	1	25	$dp[9]=dp[8]+0$	$dp[9]=5+0=5$	5 (Fits in RSU 5)
10	2	27	$dp[10]=dp[9]+1$	$dp[10]=5+1=6$	6 (New RSU needed)

Table 4: Clustering results for optimal location of RSUs

RSUs	Traffic Light Indices	Total Streams	Cluster Center Position (Centroid)
1	1, 2	5	$(1+2) \div 2 = 1.5$
2	3, 4	5	$(3+4) \div 2 = 3.5$
3	5	5	5
4	6, 7	5	$(6+7) \div 2 = 6.5$
5	8, 9	5	$(8+9) \div 2 = 8.5$
6	10	2	10

In Table 4, it describes 6 RSUs, traffic light indices covered each RSU, total video streams managed by each RSU while meeting the RSU capacity constraints, and approximate road position or centroid of each RSU by averaging traffic light indices within each cluster. The results indicate to deploy RSU 1 between traffic light index 1 and 2, RSU 2 between traffic light index 3 and 4, RSU 3 at traffic light index 5, RSU 4 between traffic light index 6 and 7, RSU 5 between traffic light index 8 and 9, and RSU 6 at traffic light index 10.

Each RSU location serves a cluster of traffic light indices with a demand sum not exceeding 5 streams, achieving optimal placement along the road. The centroids indicate where each RSU should be installed to minimize distances and ensure coverage. With 6 RSUs, this clustering balances the demand in RSUs while meeting capacity requirements.

In this work, we applied centroid based K-means clustering algorithm to determine optimal RSU locations. This algorithm is suitable for stable stream demand and equally distributed clusters. However, it is more reasonable if the video streams are uncertain and non-uniform in real world vehicular environment. In this scenario, other algorithm such as Gaussian Mixture Model (GMM) is more suitable and flexible in variably sized demand regions. Compared to K-means and GMM, K-means can allow each stream to go to one RSU while GMM can share the streams between RSUs based on the distance between vehicles and RSUs.

V. CONCLUSION AND DISCUSSION

This research proposes the framework of Cloud-based Internet of Vehicles (IoV) which integrates Roadside Units (RSUs) to enhance traffic monitoring and ensure driver safety. The incorporation of RSUs facilitates the buffering of video streams from traffic lights, leading to a significant reduction in latency when transferring data from traffic lights to cloud services. In addition, we applied dynamic programming approach for cost-effective RSU allocation and greedy algorithm to prove the optimal solution for RSU allocation obtained by dynamic programming. K-means clustering algorithm is also applied for optimal placement of RSUs. According to the results and findings, the framework not only optimizes the allocation of RSU resources but also enhances the placement of RSUs in proximity to traffic lights in order to minimize distances and ensure coverage, hence improving the efficiency of traffic data management.

For the future work, we will focus on exploring resource allocation for RSUs under the uncertainty demand of video streams from traffic lights. This investigation aims to further enhance the framework's adaptability to dynamic traffic conditions and demand fluctuation. This research is a foundational step toward advancing IoV applications and developing intelligent transportation systems that prioritize safety and efficiency.

REFERENCES

- [1] United States International Trade Commission, "Passenger vehicles: Summary of trends summary," 2013. [Online]. Available: https://www.usitc.gov/publications/332/pub_ITS_09_PassengerVehiclesSummary5211.pdf
- [2] Global Growth Insights. "Automobile Market." GLOBALGROWTH INSIGHTS.com. <https://www.globalgrowthinsights.com/market-reports/automobile-market-109995> (accessed Jun. 4, 2025).



- [3] R. Dhanare, K. K. Nagwanshi, S. Varma, and S. Pathak, "The future of internet of vehicle : challenges and applications," in *Proc. Int. Conf. Comput. Perform. Eval. (ComPE)*, Shillong, India, Dec. 2021, pp. 023–026.
- [4] F. A. Butt, J. N. Chattha, J. Ahmad, M. U. Zia, M. Rizwan, and I. H. Naqvi, "On the integration of enabling wireless technologies and sensor fusion for next-generation connected and autonomous vehicles," *IEEE Access*, vol. 10, pp. 14643–14668, 2022.
- [5] M. B. Mollah *et al.*, "Blockchain for the Internet of Vehicles towards intelligent transportation systems: A survey," *IEEE Internet Things J.*, vol. 8, no. 6, pp. 4157–4185, 2021.
- [6] R. Cherekar, "Cloud-based big data analytics: Frameworks, challenges, and future trends," *Int. J. AI, Bigdata, Comput. Manage. Stud.*, vol. 4, no. 1, pp. 35–46, 2023.
- [7] J. Beuchert, F. Solowjow, S. Trimpe, and T. Seel, "Overcoming bandwidth limitations in wireless sensor networks by exploitation of cyclic signal patterns: An event-triggered learning approach," *Sensors*, vol. 20, no. 1, 2020, Art. no. 260.
- [8] S. Akhavan Bitaghsir and A. Khonsari, "Modeling and improving the throughput of vehicular networks using cache enabled RSUs," *Telecommun. Syst.*, vol. 70, pp. 391–404, 2019.
- [9] Valuates Reports, "Global Market Share and Ranking, Overall Sales and Demand Forecast 2024-2030," 2024. [Online]. Available: <https://reports.valuates.com/market-reports/QYRE-Auto-7R15191/global-roadside-unit-rsu>
- [10] D. Liu, S. Xue, B. Zhao, B. Luo, and Q. Wei, "Adaptive dynamic programming for control: A survey and recent advances," *IEEE Trans. Syst., Man, Cybern.: Syst.*, vol. 51, no. 1, pp. 142–160, 2021.
- [11] D. Junnickel, "The greedy algorithm," in *Graphs, Networks and Algorithms*, 2nd ed. Berlin, Germany: Springer, 1999, ch. 5, pp. 129–153.
- [12] R. V. Singh and M. S. Bhatia, "Data clustering with modified K-means algorithm," in *Proc. 2011 Int. Conf. Recent Trends Inf. Technol. (ICRTIT)*, Chennai, India, Jun. 2011, pp. 717–721.
- [13] N. S. Rajput *et al.*, "Amalgamating vehicular networks with vehicular clouds, AI, and big data for next-generation ITS services," *IEEE Trans. Intell. Transp. Syst.*, vol. 25, no. 1, pp. 869–883, 2024.
- [14] Y. Sasaki *et al.*, "An edge-cloud computing model for autonomous vehicles," in *Proc. 11th Int. Workshop Planning, Perception, Navigation Intell. Vehicle*, Macau, China, Nov. 2019, pp. 99–104.
- [15] C. Liu and L. Ke, "Cloud assisted Internet of Things intelligent transportation system and the traffic control system in the smart city," *J. Control Decis.*, vol. 10, no. 2, pp. 174–187, 2023.
- [16] S. Tonayali, K. Akkaya, N. Saputro, A. S. Uluagac, and M. Nojournian, "Privacy-preserving protocols for secure and reliable data aggregation in IoT-enabled smart metering systems," *Future Gener. Comput. Syst.*, vol. 78, pp. 547–557, 2018.
- [17] X. Ma *et al.*, "QAVA: QoE-aware adaptive video bitrate aggregation for HTTP live streaming based on smart edge computing," *IEEE Trans. Broadcast.*, vol. 68, no. 3, pp. 661–676, 2022.
- [18] P. K. N. Kouonchie, V. Oduol, and G. N. Nyakoe, "Roadside units for vehicle-to-infrastructure communication: An overview," in *Proc. Sustain. Res. and Innov. Conf.*, Nairobi, Kenya, Oct. 2021, pp. 69–72.
- [19] P. Pyykonen, J. Laitinen, J. Viitanen, P. Eloranta, and T. Korhonen, "IoT for intelligent traffic system," in *IEEE 9th Int. Conf. Intell. Comput. Commun. and Process. (ICCP)*, Cluj-Napoca, Romania, Sep. 2013, pp. 175–179.

การออกแบบตัวควบคุม 2DOF-PID อย่างเหมาะสม ด้วยขั้นตอนวิธีการหาค่าเหมาะที่สุดแบบวาท

กิตติศักดิ์ ฤาแรง^{1*} ทิพา จิตหวัง² เดชา พวงดาวเรือง³

^{1*}สาขาวิชาวิศวกรรมสำรวจ คณะศิลปศาสตร์และวิทยาศาสตร์ มหาวิทยาลัยเวสเทิร์น, กาญจนบุรี, ประเทศไทย

²นักวิชาการอิสระ, นนทบุรี, ประเทศไทย

³สาขาวิชาวิศวกรรมไฟฟ้า, คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเอเชียอาคเนย์, กรุงเทพมหานคร, ประเทศไทย

*ผู้ประพันธ์บรรณกิจ อีเมล : kitiskak.l@western.ac.th

รับต้นฉบับ : 8 มีนาคม 2568; รับบทความฉบับแก้ไข : 12 กันยายน 2568; ตอบรับบทความ : 22 กันยายน 2568

เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

ตัวควบคุม PID ถูกนำเสนอครั้งแรกในปี ค.ศ.1922 และได้รับการยอมรับอย่างกว้างขวางในภาคอุตสาหกรรมมาร่วม 1 ศตวรรษ ทั้งนี้เนื่องจากตัวควบคุม PID สามารถปรับปรุงผลตอบสนองภาวะชั่วครู่และผลตอบสนองสถานะอยู่ตัว อีกทั้งยังสามารถอนุวัติได้โดยง่าย อย่างไรก็ตาม โดยธรรมชาติแล้วตัวควบคุม PID จะให้สมรรถนะของระบบมีความโดดเด่นเพียงด้านใดด้านหนึ่งตามภาวะแล็กกัน เมื่อทำการออกแบบตัวควบคุม PID ให้บรรลุวัตถุประสงค์ด้านการตามรอยสัญญาณอินพุต สมรรถนะในด้านการคุมค่าโหลดของระบบจะลดลง และในทางกลับกันก็เป็นจริง ปัญหาดังกล่าวสามารถแก้ไขได้โดยอาศัยตัวควบคุม PID แบบสองระดับขั้นความเสถียร บทความนี้นำเสนอการออกแบบตัวควบคุม 2DOF-PID อย่างเหมาะสม โดยใช้ขั้นตอนวิธีการหาค่าเหมาะที่สุดแบบวาท ซึ่งเป็นเทคนิคการค้นหาค่าเหมาะที่สุดแบบเมตาฮิวริสติกที่ทรงประสิทธิภาพ สำหรับระบบที่มีเวลาประวิงซึ่งมีผลตอบสนองล่าช้า และระบบเซอร์โวซึ่งมีผลตอบสนองรวดเร็ว ผลการออกแบบจะถูกนำไปเปรียบเทียบกับตัวควบคุม 1DOF-PID จากผลการจำลองสถานการณ์ พบว่าตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธีการหาค่าเหมาะที่สุดแบบวาท สามารถควบคุมระบบที่มีเวลาประวิงและระบบเซอร์โวได้อย่างมีประสิทธิภาพ โดยมีค่า IAE ที่ลดลงสูงสุดเท่ากับ 19.22% สำหรับระบบที่มีเวลาประวิง และ 17.14% สำหรับระบบเซอร์โว ส่งผลทำให้ผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุมค่าโหลดที่รวดเร็วและราบเรียบกว่าตัวควบคุม 1DOF-PID อย่างน่าพึงพอใจ

คำสำคัญ : ตัวควบคุม 1DOF-PID ตัวควบคุม 2DOF-PID เทคนิคการค้นหาค่าเหมาะที่สุดแบบเมตาฮิวริสติก ขั้นตอนวิธีการหาค่าเหมาะที่สุดแบบวาท

Optimal 2DOF-PID Controller Design Using Whale Optimization Algorithm

Kittisak Lurang^{1*} Thiwa Jitwang² Deacha Puangdownreong³

^{1*}*Department of Surveying Engineering, Faculty of Liberal Arts and Science, Western University,
Kanchanaburi, Thailand*

²*Independent Scholar, Nonthaburi, Thailand*

³*Department of Electrical Engineering, Faculty of Engineering, Southeast Asia University, Bangkok, Thailand*

*Corresponding Author. E-mail address: kitisak.lu@western.ac.th

Received: 8 March 2025; Revised: 12 September 2025; Accepted: 22 September 2025

Published online: 26 December 2025

Abstract

The proportional-integral-derivative (PID) controller was first introduced in 1922. It has been widely accepted in industry for almost a century, because it can improve transient and steady-state responses as well as easily implementation. However, PID controllers tend by nature to excel in one aspect of system performance due to its trade-off. When the PID controller is designed to achieve input tracking, the load regulation performance of the system is then reduced, and vice versa. This problem can be solved by using a two degree-of-freedom PID (2DOF-PID) controller. This paper presents the design of the optimal 2DOF-PID controller by using the whale optimization algorithm (WOA), one of the most efficient metaheuristic optimization techniques for the time-delayed systems having slow responses and the servo systems possessing fast responses. Results obtained by the 2DOF-PID designed by the WOA will be compared with those obtained by the 1DOF-PID controller. From the simulation results, it was found that the 2DOF-PID controller designed by the WOA algorithm can effectively control the time-delayed system and the servo system. A maximum reduction in the IAE has been achieved, with 19.22% for the time-delay system and 17.14% for the servo system. Consequently, faster and smoother tracking and load regulation responses have been satisfactorily obtained once compared to those of the 1DOF-PID controller.

Keywords: 1DOF-PID controller, 2DOF-PID controller, Metaheuristic optimization techniques, Whale optimization algorithm

1) บทนำ

ในปี ค.ศ. 1922 ตัวควบคุม PID ได้ถูกนำเสนอครั้งแรกโดย Minorsky [1] เพื่อควบคุมและรักษาเสถียรภาพของเรือรบ ตัวควบคุม PID สามารถตอบสนองวัตถุประสงค์ของการควบคุมทั้งในด้านการตามรอยสัญญาณอินพุต (input tracking) หรือการเฝ้าติดตามคำสั่ง (command following) และการคุมค่าโหลด (load regulation) หรือการกำจัดการรบกวน (disturbance rejection) ตัวควบคุม PID สามารถปรับปรุงผลตอบสนองภาวะชั่วคราวและผลตอบสนองสถานะอยู่ตัว เพื่อให้ระบบถูกควบคุมมีผลตอบสนองที่รวดเร็วและมีความแม่นยำ นอกจากนี้ การออกแบบและการอนุวัติตัวควบคุม PID สามารถกระทำได้อย่างง่ายดาย ด้วยเหตุนี้ ตัวควบคุม PID จึงได้รับการยอมรับและถูกประยุกต์ใช้ในภาคอุตสาหกรรมอย่างกว้างขวางในฐานะหัวใจหลักของระบบควบคุมอัตโนมัติมาพร้อม 1 ศตวรรษ [2], [3]

อย่างไรก็ตาม การใช้ตัวควบคุม PID เพียงตัวเดียว ซึ่งอาจเรียกว่า ตัวควบคุม PID แบบหนึ่งระดับขั้นความเสรี (one degree-of-freedom PID: 1DOF-PID controller) จะส่งผลทำให้สมรรถนะของระบบมีความโดดเด่นเพียงด้านใดด้านหนึ่ง สาเหตุที่เป็นเช่นนี้เนื่องจากธรรมชาติของระบบ 1DOF ที่มีการใช้ตัวควบคุมเพียงตัวเดียว เมื่อทำการปรับจูนตัวควบคุมให้บรรลุวัตถุประสงค์ด้านการตามรอยสัญญาณอินพุต สมรรถนะในด้านการคุมค่าโหลดของระบบจะลดลง ในทางตรงกันข้าม เมื่อทำการปรับจูนตัวควบคุมให้บรรลุวัตถุประสงค์ด้านการคุมค่าโหลด สมรรถนะในด้านการตามรอยสัญญาณอินพุตของระบบจะลดลง ปรากฏการณ์ดังกล่าวคือภาวะแลกกัน (trade-off) ของระบบ 1DOF ซึ่งการประยุกต์ในระบบอุตสาหกรรมจริง อาจส่งผลกระทบ เนื่องจากระบบควบคุมในภาคอุตสาหกรรมส่วนใหญ่มีความต้องการผลตอบสนองที่มีการตามรอยสัญญาณอินพุตและการคุมค่าโหลดที่เหมาะสมไปพร้อมกัน การควบคุมให้ระบบสามารถผลิตผลตอบสนองได้ตามวัตถุประสงค์ ทั้งในด้านการตามรอยสัญญาณอินพุตและด้านการคุมค่าโหลดไปพร้อมกัน ระบบควบคุมจะต้องมีโครงสร้างเป็น 2DOF และใช้ตัวควบคุม PID แบบสองระดับขั้นความเสรี (two degree-of-freedom PID controller: 2DOF-PID controller) [4]–[7] แนวทางการออกแบบ ตัวควบคุม 2DOF-PID แบบดั้งเดิมจะอาศัยหลักเกณฑ์การปรับจูน (tuning rules) [5], [6] ซึ่งมีข้อจำกัดเกี่ยวกับผลตอบสนองของระบบภายใต้การควบคุม (พลาเน็ต) ที่จะต้องมีลักษณะเป็นเส้นโค้งปฏิกิริยา (reaction curve) นั้นหมายความว่า ถ้าหากพลาเน็ตมีผลตอบสนองไม่เป็นไปตาม

เงื่อนไขดังกล่าว ก็จะไม่สามารถออกแบบตัวควบคุม 2DOF-PID ได้ จะเห็นได้ว่า การออกแบบตัวควบคุม 2DOF-PID อย่างเหมาะสมที่สามารถสร้างสมดุลระหว่างการตามรอยสัญญาณอินพุตและการคุมค่าโหลด สำหรับระบบควบคุมประเภทต่าง ๆ ยังคงเป็นงานที่ท้าทายในทางทฤษฎีระบบควบคุมแบบดั้งเดิม และอาจจำเป็นต้องพึ่งพาเทคนิคการหาค่าที่เหมาะสมที่สุด

ในปัจจุบัน การออกแบบตัวควบคุมได้เปลี่ยนจากแนวทางการปรับจูนแบบดั้งเดิมมาเป็นการหาค่าที่เหมาะสมที่สุดแนวใหม่โดยอาศัยเทคนิคการค้นหาค่าที่เหมาะสมที่สุดแบบเมตาฮิวริสติก (metaheuristic optimization techniques) [8] จากการสำรวจงานวิจัยที่เกี่ยวข้อง พบว่า ขั้นตอนวิธีการหาค่าที่เหมาะสมที่สุดแบบวาฬ (Whale Optimization Algorithm: WOA) เป็นเทคนิคการค้นหาค่าที่เหมาะสมที่สุดแบบเมตาฮิวริสติกที่ทรงประสิทธิภาพ ซึ่งได้รับการนำเสนอเป็นครั้งแรกโดย Mirjalili และ Lewis ในปี ค.ศ. 2016 [9] ขั้นตอนวิธี WOA ได้รับการพัฒนาขึ้นจากพฤติกรรมการล่าเหยื่อของวาฬหลังค่อม (humpback whales) ที่อาศัยเทคนิคการปล่อยเกลียวฟองอากาศเพื่อต้อนให้ฝูงปลา (เหยื่อ) มารวมกลุ่มกันบนผิวน้ำเพื่อกินเหยื่อ กลไกสำคัญของขั้นตอนวิธี WOA คือการค้นหาค่าเหยื่อ (searching for prey), การล้อมเหยื่อ (encircling prey), และการโจมตีแบบฟองอากาศ (bubble-net attacking) ขั้นตอนวิธี WOA ได้รับการทดสอบสมรรถนะกับปัญหาการหาค่าที่เหมาะสมที่สุดมาตรฐานแล้ว พบว่ามีสมรรถนะการค้นหาค่าเฉลียวกว้าง (global solution) ที่เหนือกว่าขั้นตอนวิธีการค้นหาค่าเชิงแรงโน้มถ่วง (gravitational search algorithm: GSA), วิธีวิวัฒนาการเชิงผลต่าง (Differential Evolution: DE), วิธีกลยุทธ์เชิงวิวัฒนาการ (Evolution Strategy: ES), โปรแกรมเชิงวิวัฒนาการ (Evolutionary Programming: EP), ขั้นตอนวิธีเชิงพันธุกรรม (Genetic Algorithm: GA) และการหาค่าที่เหมาะสมที่สุดแบบฝูงอนุภาค (particle swarm optimization) [9] และเนื่องจากขั้นตอนวิธี WOA เป็นเทคนิคการค้นหาแบบเมตาฮิวริสติกที่ค่อนข้างใหม่ และมีกลไกการค้นหาที่ไม่ซับซ้อน ขั้นตอนวิธี WOA จึงได้รับการประยุกต์ใช้ เพื่อแก้ปัญหาการหาค่าที่เหมาะสมที่สุดทางวิศวกรรมที่หลากหลาย อาทิเช่น การแก้ปัญหาการหาค่าที่เหมาะสมที่สุดทางวิศวกรรมไฟฟ้า [10], [12] วิศวกรรมคอมพิวเตอร์ [13], [14] วิศวกรรมอากาศยาน [15], [16] วิศวกรรมเครื่องกล [17], [18] และวิศวกรรมโยธา [19], [20] เป็นต้น ดังที่ได้รับการปริทัศน์ไว้ใน [21], [22] จากผลการสำรวจงานวิจัยที่เกี่ยวข้อง จะเห็นได้

ว่า ยังไม่มีการประยุกต์ขั้นตอนวิธี WOA เพื่อออกแบบตัวควบคุม 2DOF-PID อย่างเหมาะสม

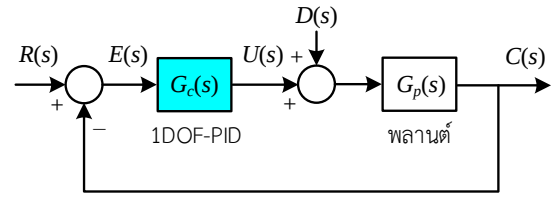
บทความวิจัยนี้ได้ประยุกต์ขั้นตอนวิธี WOA เพื่อออกแบบตัวควบคุม 2DOF-PID อย่างเหมาะสม สำหรับระบบที่มีเวลาประวิงซึ่งมีผลตอบสนองล่าช้า และระบบเซอร์โวซึ่งมีผลตอบสนองรวดเร็ว ผลการออกแบบตัวควบคุม 2DOF-PID ด้วยขั้นตอนวิธี WOA จะถูกนำไปเปรียบเทียบกับตัวควบคุม 1DOF-PID เพื่อตรวจสอบสมรรถนะในการควบคุมระบบลำดับถัดไปจะนำเสนอทฤษฎีที่เกี่ยวข้อง, วิธีดำเนินการวิจัย, ผลการวิจัยและการอภิปราย และสรุปผลการวิจัย ตามลำดับ

2) ทฤษฎีที่เกี่ยวข้อง

2.1) ระดับขั้นความเสรี

ระดับขั้นความเสรี จะพิจารณาจากจำนวนฟังก์ชันถ่ายโอนวงปิด (closed-loop transfer function) ที่สามารถปรับแต่งได้อย่างอิสระ [4]–[6] ยกตัวอย่างเช่น ระบบมีฟังก์ชันถ่ายโอนวงปิด $C(s)/R(s)$ สำหรับวัตถุประสงค์ด้านการตามรอยสัญญาณอินพุต และมีฟังก์ชันถ่ายโอนวงปิด $C(s)/D(s)$ สำหรับวัตถุประสงค์ด้านการคุมค่าโหลด ถ้าฟังก์ชันถ่ายโอนวงปิด $C(s)/R(s)$ และ $C(s)/D(s)$ สามารถปรับจูนได้อย่างอิสระต่อกัน ระบบดังกล่าวจะถูกพิจารณาเป็นระบบสองระดับขั้นความเสรี (ระบบ 2DOF) แต่ถ้าฟังก์ชันถ่ายโอนวงปิด $C(s)/R(s)$ และ $C(s)/D(s)$ ไม่สามารถปรับจูนได้อย่างอิสระต่อกัน ระบบดังกล่าวจะถูกพิจารณาเป็นระบบหนึ่งระดับขั้นความเสรี (ระบบ 1DOF)

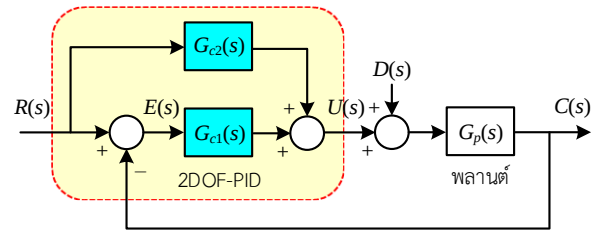
โครงสร้างของระบบ 1DOF เมื่อใช้ตัวควบคุม PID แบบดั้งเดิม ซึ่งจะพิจารณาเป็นตัวควบคุม 1DOF-PID ดังรูปที่ 1 เมื่อ $R(s)$ คือสัญญาณอินพุต, $E(s)$ คือสัญญาณความคลาดเคลื่อน, $U(s)$ คือสัญญาณควบคุม, $D(s)$ คือสัญญาณรบกวนหรือโหลด, และ $C(s)$ คือสัญญาณเอาต์พุต หรือผลตอบสนองของระบบ กำหนดให้พลานต์ (plant) มีฟังก์ชันถ่ายโอนเป็น $G_p(s)$ และตัวควบคุม 1DOF-PID มีฟังก์ชันถ่ายโอนเป็น $G_c(s)$ ดังแสดงในสมการที่ (1) เมื่อ K_c คืออัตราขยายตัวควบคุม (controller gain), τ_i คือค่าคงตัวเวลาเชิงปริพันธ์ (integral time constant), และ τ_d คือค่าคงตัวเวลาเชิงอนุพันธ์ (derivative time constant)



รูปที่ 1 : ระบบ 1DOF

$$G_c(s)|_{1\text{DOF-PID}} = K_c \left(1 + \frac{1}{\tau_i s} + \tau_d s \right) \quad (1)$$

สำหรับโครงสร้างของระบบ 2DOF เมื่อใช้ตัวควบคุม 2DOF-PID แสดงดังรูปที่ 2 และฟังก์ชันถ่ายโอนของตัวควบคุม 2DOF-PID แสดงดังสมการที่ (2) เมื่อ α และ β คือค่าตัวประกอบการลดทอน (attenuation factors) [4], [5], [6], [7]



รูปที่ 2 : ระบบ 2DOF

$$\begin{aligned} G_{c1}(s)|_{2\text{DOF-PID}} &= K_c \left(1 + \frac{1}{\tau_i s} + \tau_d s \right) \\ G_{c2}(s)|_{2\text{DOF-PID}} &= -K_c (\alpha + \beta \tau_d s) \end{aligned} \quad (2)$$

2.2) หลักเกณฑ์การปรับจูน

การออกแบบตัวควบคุม 2DOF-PID แบบดั้งเดิม จะอาศัยหลักเกณฑ์การปรับจูน AT (Araki-Taguchi tuning rules) [5], [6] ซึ่งพัฒนามาจากหลักเกณฑ์การปรับจูน ZN (Ziegler-Nichols tuning rule) วิธีที่ 1 [23], [24] โดยจะใช้ประโยชน์จากค่าเวลาประวิง L และค่าคงตัวเวลา T ที่ได้จากเส้นโค้งปฏิกิริยาของพลานต์ ซึ่งให้ผลตอบสนองต่อสัญญาณอินพุตแบบขั้นบันไดหนึ่งหน่วย (unit-step input) เป็นรูปทรงตัวอักษร S เมื่อทราบค่า L และ T จากเส้นโค้งปฏิกิริยาแล้ว การออกแบบตัวควบคุม 2DOF-PID ด้วยหลักเกณฑ์การปรับจูน AT จะอาศัยความสัมพันธ์ตามสูตรสำเร็จทางคณิตศาสตร์ ดังแสดงในตารางที่ 1

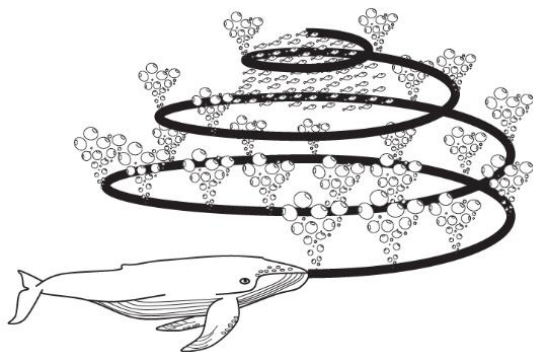
ตารางที่ 1 : หลักเกณฑ์การปรับจูน AT

Ratio L/T	2DOF-PID parameters				
	K_c	τ_i/T	τ_d/T	α	β
0.1	12.57	0.22	0.04	0.64	0.66
0.2	6.32	0.40	0.08	0.61	0.64
0.4	3.21	0.69	0.16	0.56	0.61
0.8	1.68	1.09	0.30	0.47	0.54

หลักเกณฑ์การปรับจูน AT มีข้อดีคือไม่ซับซ้อน สามารถดำเนินการได้โดยสะดวกและรวดเร็ว แต่มีข้อจำกัดคือไม่สามารถใช้กับพลานต์ที่มีผลตอบสนองไม่เป็นรูปทรงตัวอักษร S และอัตราส่วนระหว่าง L/T ต้องอยู่ในช่วง $0.1 \leq L/T \leq 0.8$ [5-6]

2.3) ขั้นตอนวิธี WOA

ขั้นตอนวิธี WOA ได้รับการพัฒนาขึ้นจากกลไกการล่าเหยื่อของวาฬหลังค่อม โดยอาศัยพฤติกรรมการปล่อยเกลียวฟองอากาศเพื่อต้อนให้ฝูงปลา (เหยื่อ) มารวมกลุ่มกัน ดังแสดง ในรูปที่ 3 วาฬหลังค่อมจะดำน้ำลงไปและเริ่มสร้างเกลียวฟองอากาศรอบตัวเหยื่อ และค่อย ๆ ว่ายน้ำขึ้นไปบนผิวน้ำเพื่อกินเหยื่อ [9]



รูปที่ 3 : การปล่อยเกลียวฟองอากาศของวาฬหลังค่อมเพื่อล่าเหยื่อ [9]

กลไกสำคัญของขั้นตอนวิธี WOA คือการค้นหาเหยื่อการล่อเหยื่อ และการโจมตีแบบฟองอากาศ สำหรับกลไกการล่อเหยื่อจะอาศัยความสัมพันธ์ทางคณิตศาสตร์ ดังแสดงในสมการที่ (3) และ (4) เมื่อ t คือจำนวนรอบการค้นหา, $\vec{X}^*$ คือตำแหน่ง (ผลเฉลย) ที่ดีที่สุด, $\vec{X}$ คือตำแหน่ง (ผลเฉลย) ปัจจุบัน, $\vec{D}$ คือเวกเตอร์ตำแหน่ง, $\vec{A}$ และ $\vec{C}$ คือเวกเตอร์สัมประสิทธิ์, $\vec{a}$ คือค่าสัมประสิทธิ์ที่ลดลงจาก $2 \rightarrow 0$ แบบเชิงเส้น ตามจำนวนรอบการค้นหา t , และ $\vec{r}_1, \vec{r}_2$ คือค่าสุ่ม (random) ที่มีการแจก

แจงแบบเอกรูป (uniform distribution) ในช่วง $[0, 1]$ เวกเตอร์สัมประสิทธิ์ $\vec{A}$ และ $\vec{C}$ ในสมการที่ (3)-(4) สามารถคำนวณได้จากสมการที่ (5) และ (6) ตามลำดับ

$$\vec{D} = |\vec{C} \cdot \vec{X}^*(t) - \vec{X}(t)| \quad (3)$$

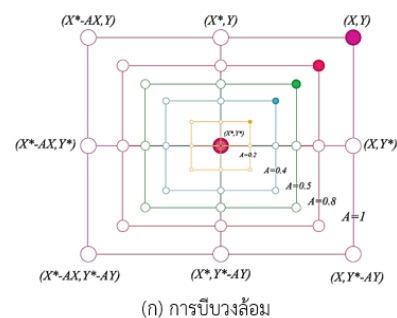
$$\vec{X}(t+1) = \vec{X}^*(t) - \vec{A} \cdot \vec{D} \quad (4)$$

$$\vec{A} = 2 \cdot \vec{a} \cdot \vec{r}_1 - \vec{a} \quad (5)$$

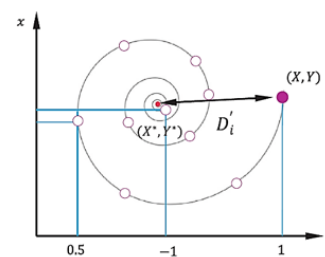
$$\vec{C} = 2 \cdot \vec{r}_2 \quad (6)$$

กลไกการโจมตีแบบฟองอากาศ จะมีกลไกที่แบ่งออกเป็น 2 ลักษณะ คือการบีบวงล้อม (shrinking encircling) และการปรับตำแหน่งแบบเกลียวกันหอย (spiral updating positions) ซึ่งวาฬหลังค่อมจะว่ายไปรอบ ๆ เหยื่อเป็นวงเกลียวที่บีบตัวไปพร้อม ๆ กัน โดยกำหนดให้มีค่าความน่าจะเป็นเท่ากับ 50%-50% ระหว่างการบีบวงล้อมและการปรับตำแหน่งแบบเกลียวกันหอย ดังแสดงในสมการที่ (7) และรูปที่ 4

$$\vec{X}(t+1) = \begin{cases} \vec{X}^*(t) - \vec{A} \cdot \vec{D} & \text{if } p < 0.5 \text{ (shrinking encircling),} \\ \vec{D}' \cdot e^{bl} \cdot \cos(2\pi l) + \vec{X}^*(t) & \text{if } p \geq 0.5 \text{ (spiral updating positions),} \end{cases} \quad (7)$$



(ก) การบีบวงล้อม



(ข) การปรับตำแหน่งแบบเกลียวกันหอย

รูปที่ 4 : กลไกการโจมตีแบบฟองอากาศ [9]

เมื่อ $\vec{D}' = |\vec{X}^*(t) - \vec{X}(t)|$, b คือค่าคงตัวสำหรับการกำหนดรูปแบบการเคลื่อนที่รูปเกลียว, และ l, p คือค่าสุ่มที่มี

การแจกแจงแบบเอกรูปในช่วง $[0, 1]$ การโจมตีแบบฟองอากาศ
เปรียบเสมือนคุณสมบัติความเข้มข้น (intensification property)
ในการค้นหาผลเฉลยเฉพาะที่ (local solution) ของขั้นตอนวิธี
WOA

กลไกการค้นหาเหยื่อ จะอาศัยความสัมพันธ์ทางคณิตศาสตร์
ดังแสดงในสมการที่ (8) และ (9)

$$\vec{D} = |\vec{C} \cdot \vec{X}_{rand}(t) - \vec{X}(t)| \quad (8)$$

$$\vec{X}(t+1) = \vec{X}_{rand}(t) - \vec{A} \cdot \vec{D} \quad (9)$$

เมื่อ $\vec{X}_{rand}(t)$ คือค่าสุ่มที่มีการแจกแจงแบบเอกรูปในช่วง
 $[0, 1]$ การค้นหาเหยื่อเปรียบเสมือนคุณสมบัติความหลากหลาย
(diversification property) ในการค้นหาผลเฉลยวงกว้าง ของ
ขั้นตอนวิธี WOA

ขั้นตอนวิธี WOA สามารถแสดงด้วยรหัสเทียม (pseudo code)
ดังแสดงในรูปที่ 5

Initialize:

- Initialize the objective function $f(x)$ and search spaces
- Initialize the whales population X_i ($i = 1, 2, \dots, n$)
- Evaluate each search agent X_i via $f(x)$
- Define X^* = the best search agent
- Define the initial iteration $t = 1$, and maximum number of iterations Max_Iter

while ($t \leq \text{Max_Iter}$)

- for** each search agent
- Update a , A , C , l , and p
- if1** ($p < 0.5$)
- if2** ($|A| < 1$)
- Update the position of the current search agent by the Eq. (3)
- else if2** ($|A| \geq 1$)
- Select a random search agent (X_{rand})
- Update the position of the current search agent by the Eq. (9)
- end if2**
- else if1** ($p \geq 0.5$)
- Update the position of the current search agent by the Eq. (7)
- end if1**
- end for**
- Check if any search agent goes beyond the search space and amend it
- Evaluate each search agent X_i via $f(x)$
- Update X^* if there is a better solution
- $t = t + 1$

end while

- Report the optimal solution X^*

รูปที่ 5 : ขั้นตอนวิธี WOA

3) วิธีดำเนินการวิจัย

3.1) พลาเน็ต

งานวิจัยนี้จะประยุกต์ขั้นตอนวิธี WOA เพื่อออกแบบตัว
ควบคุม 2DOF-PID สำหรับระบบที่มีเวลาประวิงซึ่งมีผลตอบสนอง
ล่าช้า ดังแสดงด้วยฟังก์ชันถ่ายโอน $G_{p1}(s)$ ในสมการที่ (10) [7]
และระบบเซอร์โวซึ่งมีผลตอบสนองรวดเร็ว ดังแสดงด้วยฟังก์ชัน
ถ่ายโอน $G_{p2}(s)$ ในสมการที่ (11) [25]

$$G_{p1}(s) = \frac{6}{(s+1)(s+2)(s+3)} \quad (10)$$

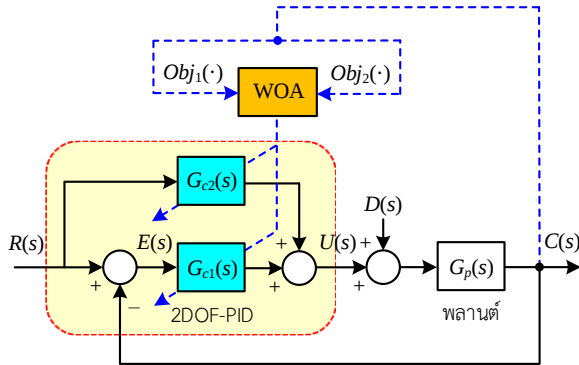
$$G_{p2}(s) = \frac{2.24}{(5.55 \times 10^{-3}s+1)(1.78 \times 10^{-3}s+1)(1.78 \times 10^{-4}s+1)} \quad (11)$$

3.2) การออกแบบตัวควบคุม 2DOF-PID ด้วยขั้นตอนวิธี WOA

การออกแบบตัวควบคุม 2DOF-PID ด้วยขั้นตอนวิธี WOA
มีการอบงานดังแสดงในรูปที่ 6 ฟังก์ชันวัตถุประสงค์ (objective
function) J แสดงดังสมการที่ (12) เมื่อ Obj_1 คือผลตอบสนอง
แบบตามรอยสัญญาณอินพุต และ Obj_2 คือผลตอบสนองแบบ
คุมค่าโหลด ค่า $0 \leq \gamma_1, \gamma_2 \leq 1$ คือตัวประกอบการปรับโทษ
(penalty factors) ซึ่งในงานวิจัยนี้กำหนดให้ $\gamma_1 = \gamma_2 = 0.5$
เนื่องจากให้ความสำคัญเท่ากัน ระหว่างผลตอบสนองแบบตาม
รอยสัญญาณอินพุตและผลตอบสนองแบบคุมค่าโหลด

จากรูปที่ 6 ฟังก์ชันวัตถุประสงค์ J จะถูกป้อนให้กับขั้นตอน
วิธี WOA เพื่อทำให้มีค่าน้อยที่สุด (minimization) ด้วยการ
ค้นหาค่าพารามิเตอร์ที่เหมาะสมของตัวควบคุม 2DOF-PID ใน
สมการที่ (2) ภายใต้ปริภูมิการค้นหาและเงื่อนไขการออกแบบ
ดังแสดงในสมการที่ (13) เมื่อ t_r คือช่วงเวลาดำเนิน (rise time), t_{r_max}
คือ t_r สูงสุดที่ยอมให้เกิดขึ้นได้, M_p คือเปอร์เซ็นต์การพุ่งเกินสูงสุด
(maximum percent overshoot), M_{p_max} คือ M_p สูงสุดที่ยอม
ให้เกิดขึ้นได้, t_s คือช่วงเวลาเข้าที่ (settling time), t_{s_max} คือ t_s
สูงสุดที่ยอมให้เกิดขึ้นได้, e_{ss} คือความคลาดเคลื่อนในสถานะ
อยู่ตัว (steady-state error), e_{ss_max} คือ e_{ss} สูงสุดที่ยอมให้เกิดขึ้น
ได้, M_{preg} คือเปอร์เซ็นต์การพุ่งเกินสูงสุดเนื่องจากการคุมค่าโหลด
(maximum percent overshoot of load regulation), M_{preg_max}
คือ M_{preg} สูงสุดที่ยอมให้เกิดขึ้นได้, t_{reg} คือช่วงเวลาคุมค่าโหลด
(load regulating time), t_{reg_max} คือ t_{reg} สูงสุดที่ยอมให้เกิดขึ้น
ได้, $[K_{c_min}, K_{c_max}]$ คือปริภูมิการค้นหาของ K_c , $[\tau_{i_min}, \tau_{i_max}]$
คือปริภูมิการค้นหาของ τ_i , $[\tau_{d_min}, \tau_{d_max}]$ คือปริภูมิการค้นหา

ของ τ_d , $[\alpha_{\min}, \alpha_{\max}]$ คือปริภูมิการค้นหาของ α , และ $[\beta_{\min}, \beta_{\max}]$ คือปริภูมิการค้นหาของ β



รูปที่ 6 : การออกแบบตัวควบคุม 2DOF-PID ด้วยขั้นตอนวิธี WOA

$$\text{Minimize } J(\cdot) = \gamma_1 \left(\frac{\text{Obj}_1(\cdot)}{\text{Obj}_{1_max}} \right) + \gamma_2 \left(\frac{\text{Obj}_2(\cdot)}{\text{Obj}_{2_max}} \right) \quad (12)$$

$$\text{Subject to } \left. \begin{aligned} &t_r \leq t_{r_max}, M_p \leq M_{p_max}, \\ &t_s \leq t_{s_max}, e_{ss} \leq e_{ss_max}, \\ &M_{preg} \leq M_{preg_max}, t_{reg} \leq t_{reg_max}, \\ &K_{c_min} \leq K_c \leq K_{c_max}, \\ &\tau_{i_min} \leq \tau_i \leq \tau_{i_max}, \\ &\tau_{d_min} \leq \tau_d \leq \tau_{d_max}, \\ &\alpha_{\min} \leq \alpha \leq \alpha_{\max}, \\ &\beta_{\min} \leq \beta \leq \beta_{\max} \end{aligned} \right\} \quad (13)$$

ในกรณีการออกแบบตัวควบคุม 1DOF-PID ด้วยขั้นตอนวิธี WOA สามารถดำเนินการได้ตามแนวทางที่กล่าวมาข้างต้น โดยการปรับค่า α และ β ในสมการที่ (13) ให้เท่ากับศูนย์ ตัวควบคุม 2DOF-PID จะกลายเป็นตัวควบคุม 1DOF-PID

3.3) การจำลองสถานการณ์และการเปรียบเทียบผล

การจำลองสถานการณ์เพื่อตรวจสอบผลตอบสนองของระบบควบคุมจะดำเนินการโดยอาศัยโปรแกรม MATLAB version 2017b (License No. #40637337) เพื่อประมวลผลบนเครื่องคอมพิวเตอร์ Intel(R) Core (TM) i7-10510U CPU@2.3 GHz, 8.0GB-RAM ขั้นตอนวิธี WOA ในรูปที่ 5 จะได้รับการพัฒนาเป็นโปรแกรมการค้นหาด้วยโปรแกรม MATLAB เพื่อออกแบบตัวควบคุม 2DOF-PID

สำหรับระบบที่มีเวลาประวิง $G_{p1}(s)$ ในสมการที่ (10) ซึ่งมีผลตอบสนองที่เป็นเส้นโค้งปฏิกิริยาในลักษณะเป็นรูปทรงตัวอักษร S ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA จะถูก

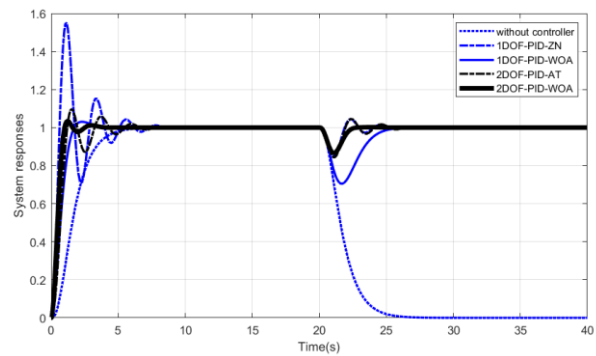
นำไปเปรียบเทียบกับตัวควบคุม 2DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน AT, ตัวควบคุม 1DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน ZN, และตัวควบคุม 1DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA

สำหรับระบบเซอโรโว $G_{p2}(s)$ ในสมการที่ (11) ซึ่งมีผลตอบสนองไม่เป็นรูปทรงตัวอักษร S ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA จะถูกเปรียบเทียบกับตัวควบคุม 1DOF-PID ที่ออกแบบด้วยวิธี Modulus optimum (MO) [25] และตัวควบคุม 1DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA

4) ผลการวิจัยและการอภิปราย

4.1) ผลการออกแบบตัวควบคุมสำหรับระบบที่มีเวลาประวิง

ระบบที่มีเวลาประวิง $G_{p1}(s)$ ในสมการที่ (10) ให้ผลตอบสนองเป็นเส้นโค้งปฏิกิริยาในลักษณะเป็นรูปทรงตัวอักษร S ดังแสดงในรูปที่ 7 โดยมีค่า $L = 0.4211$ วินาที และ $T = 2.3157$ วินาที



รูปที่ 7 : ผลตอบสนองของระบบที่มีเวลาประวิง $G_{p1}(s)$

การออกแบบตัวควบคุม 1DOF-PID สำหรับพลาเน็ต $G_{p1}(s)$ ด้วยหลักเกณฑ์การปรับจูน ZN [12-13] อาศัยความสัมพันธ์ $K_c = 1.2T/L = 6.599$, $\tau_i = 2L = 0.8422$ วินาที, และ $\tau_d = 0.5L = 0.2105$ วินาที ตัวควบคุม 1DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน ZN แสดงดังสมการที่ (14) และให้ผลตอบสนอง ดังแสดงในรูปที่ 7

$$G_c(s)|_{\text{1DOF-PID-ZN}} = 6.60 + \frac{7.84}{s} + 1.39s \quad (14)$$

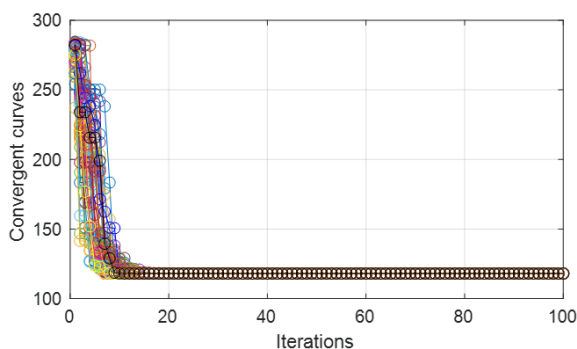
การออกแบบตัวควบคุม 2DOF-PID สำหรับพลาเน็ต $G_{p1}(s)$ ด้วยหลักเกณฑ์การปรับจูน AT [5-6] จะพบว่า อัตราส่วน $L/T = 0.1818 \approx 0.2$ ดังนั้น จากตารางที่ 1 จะได้ $K_c = 6.32$, $\tau_i = 0.4T = 0.9263$ วินาที, $\tau_d = 0.08T = 0.1853$ วินาที, $\alpha = 0.61$

และ $\beta = 0.64$ ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน AT แสดงดังสมการที่ (15) และให้ผลตอบสนอง ดังแสดงในรูปที่ 7

$$\left. \begin{aligned} G_{c1}(s)|_{2DOF-PID-AT} &= 6.32 + \frac{6.82}{s} + 1.17s \\ G_{c2}(s)|_{2DOF-PID-AT} &= -(3.86 + 0.75s) \end{aligned} \right\} \quad (15)$$

การออกแบบตัวควบคุม 1DOF-PID และ 2DOF-PID สำหรับพลาเน็ต $G_{p1}(s)$ ด้วยขั้นตอนวิธี WOA เริ่มต้นจากการทดสอบเพื่อหาจำนวนวาทที่เหมาะสม โดยการปรับจำนวนวาท $n = 5, 10, \dots, 50$ ซึ่งพบว่า $n = 30$ คือค่าที่เหมาะสมสำหรับการออกแบบ กำหนดให้เกณฑ์ยุติการค้นหาค้นหา (Termination Criteria: TC) คือจำนวนรอบการค้นหามากที่สุด $Max_Iter = 100$ เนื่องจากเป็นจำนวนรอบสูงสุดที่สามารถรับประกันได้ว่าการค้นหาค้นหาสามารถลู่เข้าหาผลเฉลยได้อย่างแน่นอน และดำเนินการค้นหาทั้งหมด 50 ครั้ง (trials) เพื่อออกแบบตัวควบคุม 1DOF-PID และ 2DOF-PID ที่เหมาะสม อันจะส่งผลทำให้ฟังก์ชันวัตถุประสงค์ J ในสมการที่ (12) มีค่าน้อยที่สุด โดยกำหนดให้ปริภูมิการค้นหาค้นหาและเงื่อนไขการออกแบบในสมการที่ (13) แสดงดังสมการที่ (16)

$$\left. \begin{aligned} \text{Subject to } t_r &\leq 2.00 \text{ s}, M_p \leq 15\%, \\ t_s &\leq 7.50 \text{ s}, e_{ss} \leq 0.01\%, \\ M_{preg} &\leq 30\%, t_{reg} \leq 5.00 \text{ s}, \\ 0 &\leq K_c \leq 10.00, \\ 0 &\leq \tau_i \leq 5.00 \text{ s}, \\ 0 &\leq \tau_d \leq 5.00 \text{ s}, \\ 0 &\leq \alpha \leq 2.50, \\ 0 &\leq \beta \leq 2.50 \end{aligned} \right\} \quad (16)$$



รูปที่ 8 : เส้นโค้งการลู่เข้าของฟังก์ชันวัตถุประสงค์ในการออกแบบตัวควบคุม 2DOF-PID สำหรับพลาเน็ต $G_{p1}(s)$ ด้วยขั้นตอนวิธี WOA

รูปที่ 8 แสดงเส้นโค้งการลู่เข้าของฟังก์ชันวัตถุประสงค์ J ในการออกแบบตัวควบคุม 2DOF-PID สำหรับพลาเน็ต $G_{p1}(s)$

ด้วยขั้นตอนวิธี WOA จากการดำเนินการค้นหาทั้งหมด 50 ครั้ง จากรูปที่ 8 จะพบว่า ขั้นตอนวิธี WOA สามารถออกแบบตัวควบคุม 2DOF-PID ได้อย่างมีประสิทธิภาพ ซึ่งมีค่าเฉลี่ยเท่ากับ 118.02 และค่าเบี่ยงเบนมาตรฐานเท่ากับ 1.26×10^{-6} สำหรับเส้นโค้งการลู่เข้าของฟังก์ชันวัตถุประสงค์ J ในการออกแบบตัวควบคุม 1DOF-PID มีค่าเฉลี่ยเท่ากับ 124.73 และค่าเบี่ยงเบนมาตรฐานเท่ากับ 3.58×10^{-5} และจะมีลักษณะเช่นเดียวกับรูปที่ 8 จึงไม่น่าสนใจในที่นี้ ขั้นตอนวิธี WOA ใช้เวลาเฉลี่ยในการออกแบบตัวควบคุม 1DOF-PID และ 2DOF-PID สำหรับพลาเน็ต $G_{p1}(s)$ เท่ากับ 7.12 วินาที และ 9.46 วินาที ตัวควบคุม 1DOF-PID และ 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA สำหรับพลาเน็ต $G_{p1}(s)$ แสดงดังสมการที่ (17) และ (18) และให้ผลตอบสนองดังรูปที่ 7

$$G_c(s)|_{1DOF-PID-WOA} = 2.11 + \frac{1.45}{s} + 0.76s \quad (17)$$

$$\left. \begin{aligned} G_{c1}(s)|_{2DOF-PID-WOA} &= 6.24 + \frac{6.28}{s} + 2.13s \\ G_{c2}(s)|_{2DOF-PID-WOA} &= -(1.98 + 0.83s) \end{aligned} \right\} \quad (18)$$

ตารางที่ 2 : การเปรียบเทียบผลตอบสนองของระบบที่มีเวลาประวิง

ตัวควบคุม	ผลตอบสนอง								
	ตามรอยสัญญาณอินพุต					คุ่มค่าโหลด			IAE_Total
	t_r (s)	M_p (%)	t_s (s)	e_{ss} (%)	IAE1	M_{preg} (%)	t_{reg} (s)	IAE2	
ไม่มีตัวควบคุม	3.37	0.00	5.00	0.00	58.51	ไม่สามารถคุ่มค่าได้			
1DOF-PID-ZN	0.62	55.04	6.83	0/00	49.39	14.65	3.69	1.87	51.26
1DOF-PID-WOA	1.72	3.05	2.94	0.00	47.81	29.50	4.73	5.05	52.86
2DOF-PID-AT	1.20	9.70	5.21	0.00	48.87	15.71	3.83	2.06	50.93
2DOF-PID-WOA	1.01	3.42	2.11	0.00	41.11	13.48	2.17	1.59	42.70

จากรูปที่ 7 ผลตอบสนองแบบตามรอยสัญญาณอินพุตแสดงในช่วงเวลา 0–20 วินาที และผลตอบสนองแบบคุ่มค่าโหลดแสดงในช่วงเวลา 20–40 วินาที การเปรียบเทียบผลตอบสนองของระบบที่มีเวลาประวิง $G_{p1}(s)$ แสดงดังตารางที่ 2 ซึ่งพบว่า ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA สามารถควบคุมระบบ $G_{p1}(s)$ ได้อย่างมีประสิทธิภาพ สอดคล้องกับเงื่อนไข

การออกแบบที่กำหนดในสมการที่ (16) และให้ผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลดที่รวดเร็วและราบเรียบกว่าตัวควบคุม 2DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน AT, ตัวควบคุม 1DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน ZN, และตัวควบคุม 1DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA อย่างน่าพึงพอใจ

สมรรถนะด้านการตามรอยสัญญาณอินพุตและการคุ่มค่าโหลดของตัวควบคุม 1DOF-PID-ZN, 1DOF-PID-WOA, 2DOF-PID-AT และ 2DOF-PID-WOA สำหรับระบบที่มีเวลาประวิง ในรูปที่ 7 สามารถแสดงด้วยค่าปริพันธ์ของค่าความคลาดเคลื่อนสัมบูรณ์ (Integral of Absolute Error: IAE) ระหว่าง สัญญาณอินพุตอ้างอิงและสัญญาณเอาต์พุตของระบบ ดังแสดงในตารางที่ 2 เมื่อ IAE1 คือค่า IAE ของการตามรอยสัญญาณอินพุต, IAE2 คือค่า IAE ของการคุ่มค่าโหลด และ IAE_Total คือผลรวมของค่า IAE1 และ IAE2 ซึ่งพบว่า ตัวควบคุม 2DOF-PID-WOA สามารถลดค่า IAE_Total ลงได้ 19.22%, 16.70% และ 16.16% เมื่อเปรียบเทียบกับตัวควบคุม 1DOF-PID-WOA, 1DOF-PID-ZN และ 2DOF-PID-AT ตามลำดับ แสดงให้เห็นว่า ตัวควบคุม 2DOF-PID-WOA สามารถผลิตผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลดได้อย่างเหมาะสมและสมดุล ทำให้ผลตอบสนองในภาพรวมมีความรวดเร็วและราบเรียบกว่าตัวควบคุม 1DOF-PID-WOA, 1DOF-PID-ZN และ 2DOF-PID-AT ตามลำดับ

4.2) ผลการออกแบบตัวควบคุมสำหรับระบบเซอร์โว

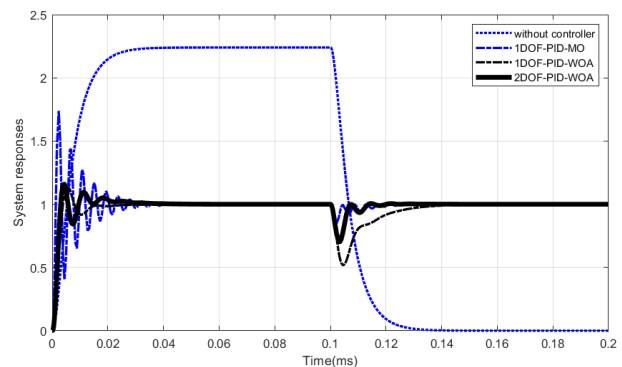
ระบบเซอร์โว $G_{p2}(s)$ ในสมการที่ (11) ให้ผลตอบสนองไม่เป็นรูปทรงตัวอักษร S ดังแสดงในรูปที่ 9 ตัวควบคุม 1DOF-PID ที่

ออกแบบด้วยวิธี MO [14] สำหรับ พลานต์ $G_{p2}(s)$ แสดงดังสมการที่ (19)

$$G_c(s)|_{1DOF-PID-MO} = 9.20 + \frac{1,256.22}{s} + 0.01s \quad (19)$$

การออกแบบตัวควบคุม 1DOF-PID และ 2DOF-PID สำหรับพลานต์ $G_{p2}(s)$ ด้วยขั้นตอนวิธี WOA กำหนดให้จำนวนวาท $n = 30$ TC คือ $Max_Iter = 100$ และดำเนินการค้นหาทั้งหมด 50 ครั้ง เพื่อออกแบบตัวควบคุม 1DOF-PID และ 2DOF-PID ที่เหมาะสม อันจะส่งผลทำให้ฟังก์ชันวัตถุประสงค์ J ในสมการที่ (12) มีค่าน้อยที่สุด โดยกำหนดให้ปริภูมิการค้นหาและเงื่อนไขการออกแบบในสมการที่ (13) แสดงดังสมการที่ (20)

$$\left. \begin{aligned} \text{Subject to } & t_r \leq 10.00 \mu s, M_p \leq 20\%, \\ & t_s \leq 0.50 \text{ ms}, e_{ss} \leq 0.01\%, \\ & M_{preg} \leq 50\%, t_{reg} \leq 0.75 \text{ ms}, \\ & 0 \leq K_c \leq 10.00, \\ & 0 \leq \tau_i \leq 5.00 \text{ s}, \\ & 0 \leq \tau_d \leq 5.00 \text{ s}, \\ & 0 \leq \alpha \leq 2.50, \\ & 0 \leq \beta \leq 2.50 \end{aligned} \right\} \quad (20)$$



รูปที่ 9 : ผลตอบสนองของระบบเซอร์โว $G_{p2}(s)$

ตารางที่ 3 : การเปรียบเทียบผลตอบสนองของระบบเซอร์โว

ตัวควบคุม	ผลตอบสนอง								
	ตามรอยสัญญาณอินพุต					คุมค่าโหลต			IAE_Total
	t_r (ms)	M_p (%)	t_s (ms)	e_{ss} (%)	IAE1	M_{preg} (%)	t_{reg} (ms)	IAE2	
ไม่มีตัวควบคุม	0.015	0.00	0.024	124.0	146.77	ไม่สามารถคุมค่าได้			
1DOF-PID-MO	1.29×10^{-3}	73.85	0.032	0.00	19.18	15.15	0.016	1.18	20.36
1DOF-PID-WOA	3.85×10^{-3}	14.27	0.014	0.00	15.65	48.24	0.034	6.41	22.06
2DOF-PID-WOA	3.07×10^{-3}	15.98	0.026	0.00	15.96	29.72	0.012	2.32	18.28

เส้นโค้งการลู่เข้าของฟังก์ชันวัตถุประสงค์ J ในการออกแบบตัวควบคุม 1DOF-PID สำหรับพลาเน็ต $G_{p2}(s)$ ด้วยขั้นตอนวิธี WOA จะมีค่าเฉลี่ยเท่ากับ 136.67 และค่าเบี่ยงเบนมาตรฐานเท่ากับ 4.06×10^{-5} และสำหรับตัวควบคุม 2DOF-PID จะมีค่าเฉลี่ยเท่ากับ 122.58 และค่าเบี่ยงเบนมาตรฐานเท่ากับ 2.02×10^{-6} ซึ่งจะมีลักษณะเป็นเส้นโค้งเช่นเดียวกับรูปที่ 8 จึงไม่ขอนำเสนอในที่นี้

ขั้นตอนวิธี WOA ใช้เวลาเฉลี่ยในการออกแบบตัวควบคุม 1DOF-PID และ 2DOF-PID สำหรับพลาเน็ต $G_{p2}(s)$ เท่ากับ 7.14 วินาที และ 9.52 วินาที ตัวควบคุม 1DOF-PID และ 2DOF-PID ที่ออกแบบด้วยด้วยขั้นตอนวิธี WOA สำหรับพลาเน็ต $G_{p2}(s)$ แสดงดังสมการที่ (21) และ (22) และให้ผลตอบสนองดังรูปที่ 9

$$G_c(s)|_{1DOF-PID-WOA} = 1.74 + \frac{192.85}{s} + 0.001s \quad (21)$$

$$\left. \begin{aligned} G_{c1}(s)|_{2DOF-PID-WOA} &= 2.67 + \frac{796.45}{s} + 0.01s \\ G_{c2}(s)|_{2DOF-PID-WOA} &= -(0.83 + 0.02s) \end{aligned} \right\} \quad (22)$$

จากรูปที่ 9 ผลตอบสนองแบบตามรอยสัญญาณอินพุตแสดงในช่วงเวลา 0–0.1 มิลลิวินาที และผลตอบสนองแบบคุ่มค่าโหลตแสดงในช่วงเวลา 0.1–0.2 มิลลิวินาที การเปรียบเทียบผลตอบสนองของระบบที่มีเวลาประวิง $G_{p2}(s)$ ดังตารางที่ 3 ซึ่งพบว่า ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA สามารถควบคุมระบบ $G_{p2}(s)$ ได้อย่างมีประสิทธิภาพ สอดคล้องกับเงื่อนไขการออกแบบที่กำหนดในสมการที่ (20) และให้ผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลตที่รวดเร็วและราบเรียบกว่าตัวควบคุม 1DOF-PID ที่ออกแบบด้วยวิธี MO และขั้นตอนวิธี WOA อย่างน่าพึงพอใจ

สมรรถนะด้านการตามรอยสัญญาณอินพุตและการคุ่มค่าโหลต ของตัวควบคุม 1DOF-PID-MO, 1DOF-PID-WOA และ 2DOF-PID-WOA สำหรับระบบเซอร์โว ในรูปที่ 9 สามารถแสดงด้วยค่า IAE ระหว่างสัญญาณอินพุตอ้างอิงและสัญญาณเอาต์พุตของระบบ ดังแสดงในตารางที่ 3 เมื่อ IAE1 คือค่า IAE ของการตามรอยสัญญาณอินพุต, IAE2 คือค่า IAE ของการคุ่มค่าโหลต และ IAE_Total คือผลรวมของค่า IAE1 และ IAE2 ซึ่งพบว่า ตัวควบคุม 2DOF-PID-WOA สามารถลดค่า IAE_Total ลงได้ 17.14% และ 10.22% เมื่อเปรียบเทียบกับตัวควบคุม 1DOF-PID-WOA และ 1DOF-PID-MO ตามลำดับ นั้นหมายความว่า ตัวควบคุม 2DOF-PID-WOA สามารถผลิตผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลตได้อย่างเหมาะสมและสมคูล ทำให้ผลตอบสนองในภาพรวมมีความรวดเร็ว

และราบเรียบกว่าตัวควบคุม 1DOF-PID-WOA และ 1DOF-PID-MO ตามลำดับ

5) สรุป

บทความนี้นำเสนอการออกแบบตัวควบคุม 2DOF-PID อย่างเหมาะสม สำหรับระบบที่มีเวลาประวิงและระบบเซอร์โว โดยใช้ขั้นตอนวิธีการหาค่าเหมาะที่สุดแบบวาท (WOA) จากผลการออกแบบและการจำลองสถานการณ์ พบว่า ขั้นตอนวิธี WOA สามารถออกแบบตัวควบคุม 2DOF-PID ได้อย่างมีประสิทธิภาพ สอดคล้องกับเงื่อนไขการออกแบบที่กำหนด ในกรณีของระบบที่มีเวลาประวิง พบว่า ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA ให้ผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลตที่รวดเร็วและราบเรียบกว่าตัวควบคุม 2DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน AT, ตัวควบคุม 1DOF-PID ที่ออกแบบด้วยหลักเกณฑ์การปรับจูน ZN, และตัวควบคุม 1DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA และในกรณีของระบบเซอร์โว พบว่า ตัวควบคุม 2DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA ให้ผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลต ที่รวดเร็วและราบเรียบกว่าตัวควบคุม 1DOF-PID ที่ออกแบบด้วยวิธี MO และขั้นตอนวิธี WOA เมื่อพิจารณาจากผลการออกแบบในภาพรวมของงานวิจัยนี้ พบว่า ตัวควบคุม 2DOF-PID สามารถควบคุมให้ระบบมีผลตอบสนองที่น่าพึงพอใจทั้งผลตอบสนองแบบตามรอยสัญญาณอินพุตและผลตอบสนองแบบคุ่มค่าโหลต ในขณะที่ตัวควบคุม 1DOF-PID สามารถควบคุมให้ระบบมีผลตอบสนองที่น่าพึงพอใจเฉพาะผลตอบสนองแบบตามรอยสัญญาณอินพุตเท่านั้น แต่จะให้ผลตอบสนองแบบคุ่มค่าโหลตที่มีคุณภาพต่ำจากการวิเคราะห์ผลในภาพรวม การที่ตัวควบคุม 2DOF-PID สามารถผลิตผลตอบสนองได้ตามวัตถุประสงค์ ทั้งในด้านการตามรอยสัญญาณอินพุตและด้านการคุ่มค่าโหลตไปพร้อมกันได้ดีกว่าตัวควบคุม 1DOF-PID เนื่องจากตัวควบคุม 2DOF-PID มีส่วนของตัวควบคุม $G_{c1}(s)$ และ $G_{c2}(s)$ ดังแสดงในสมการที่ (2) ที่สามารถแยกการควบคุมได้อย่างอิสระ เมื่อนำขั้นตอนวิธี WOA ซึ่งเป็นเทคนิคการค้นหาค่าเหมาะที่สุดแบบเมตาฮิวริสติกที่ทรงประสิทธิภาพ มาช่วยในการออกแบบค่าพารามิเตอร์ที่เหมาะสม จึงทำให้ตัวควบคุม 2DOF-PID สามารถให้ผลตอบสนองที่เหนือกว่าตัวควบคุม 2DOF-PID และตัวควบคุม 1DOF-PID ที่ออกแบบ

ด้วยหลักเกณฑ์ การปรับปรุง และให้ผลตอบสนองที่เหนือกว่าตัวควบคุม 1DOF-PID ที่ออกแบบด้วยขั้นตอนวิธี WOA

แนวทางการดำเนินงานวิจัยในอนาคต สามารถแบ่งออกได้หลายแนวทาง ดังนี้ แนวทางแรก คือการนำผลการออกแบบตัวควบคุมในบทความนี้ไปอนุวัติ (implementation) เพื่อนำผลการทดสอบจริงมาเปรียบเทียบกับผลการจำลองสถานการณ์แนวทางต่อมา คือการพัฒนาขั้นตอนวิธี WOA ให้มีประสิทธิภาพในการค้นหาผลเฉลยวงกว้างที่รวดเร็วยิ่งขึ้น และแนวทางสุดท้ายคือการประยุกต์ขั้นตอนวิธี WOA เพื่อการระบุเอกลักษณ์ระบบ (system identification) และออกแบบตัวควบคุมให้กับระบบจริงที่มีความซับซ้อนและหลากหลายมากยิ่งขึ้น

REFERENCES

- [1] N. Minorsky, "Directional stability of automatically steered bodies," *J. Amer. Soc. Naval Eng.*, vol. 34, no. 2, pp. 280–309, May 1922.
- [2] K. J. Åström and T. Häggglund, *PID Controllers: Theory, Design, and Tuning*, 2nd ed. North Carolina, NC, USA: Instrument Society of America, 1995.
- [3] A. O'Dwyer, *Handbook of PI and PID Controller Tuning Rules*, London, UK: Imperial College Press, 2003.
- [4] M. Araki, "Two-degree-of-freedom control system: part I," *Syst. Control*, vol. 29, pp. 649–656, 1985.
- [5] M. Araki and H. Taguchi, "Two-degree-of-freedom PID controllers," *Syst. Control Inf.*, vol. 42, pp. 18–25, 1998.
- [6] H. Taguchi and M. Araki, "Two-degree-of-freedom PID controllers: their functions and optimal tuning," in *Proc. IFAC Workshop in Digit. Control 2000: Past, Present and Future of PID Control*, Apr. 2000, pp. 91–96.
- [7] K. Lurang and D. Puangdownreong, "Two-degree-of-freedom PIDA controllers design optimization for liquid-level system by using modified bat algorithm," *Int. J. Innov. Comput. Inf. Control*, vol. 16, no. 2, pp. 715–732, Apr. 2020.
- [8] V. Zakian, *Control Systems Design: A New Framework*, London, UK: Springer-Verlag London, 2005.
- [9] S. Mirjalili and A. Lewis, "The whale optimization algorithm," *Adv. Eng. Softw.*, vol. 95, pp. 51–67, May 2016.
- [10] D. Prakash and C. Lakshminarayana, "Optimal siting of capacitors in radial distribution network using whale optimization algorithm," *Alexandria Eng. J.*, vol. 56, no. 4, pp. 499–509, Dec. 2017.
- [11] H. J. Touma, "Study of the economic dispatch problem on IEEE 30-bus system using whale optimization algorithm," *Int. J. Eng., Sci. Technol.*, vol. 5, no. 1, pp. 11–18, Jun. 2016.
- [12] O. Baimakhanov, A. Saukhimov, and O. Ceylan, "Whale optimization algorithm based optimal operation of power distribution systems," in *Proc. 57th Int. Univ. Power Eng. Conf. (UPEC)*, 2022, pp. 1–6, doi: 10.1109/UPEC55022.2022.9917748.
- [13] A. N. Jadhav and N. Gomathi, "WGC: hybridization of exponential grey wolf optimizer with whale optimization for data clustering," *Alexandria Eng. J.*, vol. 57, no. 3, pp. 1569–1584, Sep. 2018.
- [14] U. Dixit, A. Mishra, A. Shukla, and R. Tiwari, "Texture classification using convolutional neural network optimized with whale optimization algorithm," *SN Appl. Sci.*, vol. 1, no. 655, pp. 1–11, May 2019.
- [15] S. Aftab, N. Razak, A. S. M. Rafie, and K. A. Ahmad, "Mimicking the humpback whale: An aerodynamic perspective," *Prog. Aerosp. Sci.*, vol. 84, pp. 48–69, Jul. 2016.
- [16] X. Huang, R. Wang, X. Zhao, and K. Hu, "Aero-engine performance optimization based on whale optimization algorithm," in *Proc. 36th Chinese Control Conf. (CCC)*, 2017, pp. 11437–11441, doi: 10.23919/ChiCC.2017.8029182.
- [17] A. H. Kashan, "An efficient algorithm for constrained global optimization and application to mechanical engineering design: league championship algorithm (LCA)," *Comput.-Aided Des.*, vol. 43, no. 12, pp. 1769–1792, Dec. 2011.
- [18] G. Hou, L. Gong, Z. Yang, and J. Zhang, "Multi-objective economic model predictive control for gas turbine system based on quantum simultaneous whale optimization algorithm," *Energy Convers. Manage.*, vol. 207, Mar. 2020, Art. no. 112498.
- [19] A. Kaveh and M. I. Ghazaan, "Enhanced whale optimization algorithm for sizing optimization of skeletal structures," *Mech. Based Des. Struct. Mach.*, vol. 45, no. 3, pp. 345–362, Sep. 2016.
- [20] M. Rohani, G. Shafabakhsh, A. Haddad, and E. Asnaashari, "The workflow planning of construction sites using whale optimization algorithm (WOA)," *Turkish Online J. Des. Art Commun.*, vol. 6, pp. 2938–2950, Nov. 2016.
- [21] N. Rana, M. S. A. Latiff, S. M. Abdulhamid, and H. Chiroma, "Whale optimization algorithm: a systematic review of contemporary applications, modifications and developments," *Neural Comput. Appl.*, vol. 32, pp. 16245–16277, Mar. 2020.



- [22] H. N.-S. Mohammad, H. Zamani, Z. A. Varzaneh, and S. Mirjalili, "A systematic review of the whale optimization algorithm: theoretical foundation, improvements, and hybridizations," *Arch. Comput. Methods Eng.*, vol. 30, pp. 4113–4159, May 2023.
- [23] J. G. Ziegler and N. B. Nichols, "Optimum settings for automatic controllers," *Trans. ASME*, vol. 64, no. 8, pp. 759–765, Nov. 1942.
- [24] J. G. Ziegler and N. B. Nichols, "Process lags in automatic control circuits," *Trans. ASME*, vol. 65, no. 5, pp. 433–444, Jul. 1943.
- [25] D. Puangdownreong, Y. Prempraneerat, and S. Sujitjorn, "Development of a dc-servo motor control system for testing of the Eitelberg's method," (in Thai), in *Proc. 18th Elect. Eng. Conf. (EECON-18)*, pp. 809–814, 1995.

ผลกระทบของฝนต่อพฤติกรรมของดิน : การประเมินความเสี่ยงดินถล่ม ด้วยแบบจำลองทางกายภาพ

ลัดดาวัลย์ ดุลย์^{1*} ทวีชัย ยี่เมา² ศุภกฤต กันทาคำ³

^{1*,2,3}สาขาวิศวกรรมเครื่องกล คณะวิศวกรรมศาสตร์ มหาวิทยาลัยเทคโนโลยีราชมงคลล้านนา, เชียงใหม่, ประเทศไทย

*อีเมลผู้ประพันธ์บรรณกิจ : laddawon_dul@rmutl.ac.th

รับต้นฉบับ : 13 มกราคม 2568; รับบทความฉบับแก้ไข : 7 พฤษภาคม 2568; ตอบรับบทความ : 27 พฤษภาคม 2568

เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

ดินถล่มถือเป็นหนึ่งในภัยพิบัติทางธรณีวิทยาที่สร้างความเสียหายรุนแรงในพื้นที่ภูเขา ซึ่งมักมีปัจจัยเสี่ยงจากปริมาณน้ำฝน คุณสมบัติของดินและการตัดถนนผ่านในพื้นที่เขาสูง จึงทำการศึกษาและวิเคราะห์พฤติกรรมของดินภายใต้เงื่อนไขฝนตกต่อเนื่องในพื้นที่เสี่ยง บริเวณทางหลวงหมายเลข 1096 อำเภอสะเมิง จังหวัดเชียงใหม่ โดยใช้การจำลองทางกายภาพในห้องปฏิบัติการ เพื่อเปรียบเทียบพฤติกรรมของมวลดิน 2 ประเภท ได้แก่ มวลดินเดิมบริเวณข้างทาง มีหน่วยน้ำหนัก 15.30 kN/m^3 และมวลดินบดอัดที่เป็นไหลทาง มีหน่วยน้ำหนัก 18.34 kN/m^3 ภายใต้ปริมาณน้ำฝน 4 ระดับ ผลการทดลองพบว่า มวลดินเดิมเกิดการถล่มเมื่อมีฝนตกที่อัตรา 60 มิลลิเมตรต่อชั่วโมง ขณะที่มวลดินบดอัดไหลทางถล่มเมื่อปริมาณฝนถึง 160 มิลลิเมตรต่อชั่วโมง โดยค่าดัชนีอัตราการถล่มต่อปริมาณน้ำฝนที่ 0.75 เป็นจุดขีดจำกัดเริ่มเกิดการถล่มอย่างต่อเนื่องในมวลดินเดิม และเกิดการถล่มแบบเฉียบพลันในมวลดินบดอัดเมื่อฝนเกิน 220 มิลลิเมตรต่อชั่วโมง จากผลการศึกษาได้เสนอใช้เกณฑ์ปริมาณน้ำฝนร่วมกับดัชนีอัตราการถล่มเป็นเครื่องมือเตือนภัยเบื้องต้น โดยกำหนดระดับเตือนภัยสีเหลือง ที่ช่วง 60–160 มิลลิเมตรต่อชั่วโมง ที่ควรเริ่มเฝ้าระวังการเคลื่อนตัวของดิน ระดับสีเหลืองเข้มเพื่อเตรียมการอพยพ และระดับสีแดงควรปิดถนนเมื่อฝนตกเกิน 220 มิลลิเมตรต่อชั่วโมง นอกจากนี้ ยังสามารถจำแนกพฤติกรรมการถล่มได้ 4 ลักษณะ ซึ่งเป็นข้อมูลสำคัญต่อการวางแผนป้องกันและบริหารจัดการความเสี่ยงในพื้นที่เสี่ยงดินถล่มอย่างมีประสิทธิภาพ

คำสำคัญ : ดัชนีเตือนภัยล่วงหน้า ดินถล่ม แบบจำลองทางกายภาพ ปริมาณน้ำฝน เสถียรภาพของลาดชันของดิน

Rainfall Impact on Soil Behavior: Landslide Risk Assessment Using Physical Modeling

Laddawon Dul^{1*} Taweechai Yeemao² Suphakrit Kantakam³

^{1*,2,3}*Mechanical Engineering Department, Faculty of Engineering, Rajamangala University of Technology Lanna,
Chiang Mai, Thailand*

*Corresponding Author. E-mail address: laddawon_dul@rmutl.ac.th

Received: 3 January 2025; Revised: 7 May 2025; Accepted: 27 May 2025

Published online: 26 December 2025

Abstract

Landslides are recognized as major geological hazards that can cause significant damage in mountainous areas, where risk factors often include high rainfall, specific soil characteristics, and road construction across elevated terrain. This study considers to analyze soil behavior under continuous rainfall conditions in a high-risk area along Highway No. 1096, Samoeng District, Chiang Mai Province, using physical modeling in a laboratory setting. Two types of soil masses were examined: natural roadside soil with a unit weight of 15.30 kN/m³ and compacted shoulder soil with a unit weight of 18.34 kN/m³, tested under four rainfall intensity levels. The results revealed that the natural soil mass collapsed when rainfall reached 60 mm/hr, while the compacted shoulder soil failed at 160 mm/hr. A landslide intensity index of 0.75 was identified as a critical threshold, marking the onset of continuous failures in natural soil and sudden failure in compacted soil when rainfall exceeded 220 mm/hr. Based on these findings, the study proposes a preliminary warning system using both rainfall thresholds and the landslide intensity index. A Yellow Alert level 60–160 mm/hr is recommended for initial monitoring of soil movement, Dark Yellow for evacuation preparedness, and Orange Alert for road closure when rainfall exceeds 220 mm/hr. Additionally, four distinct types of landslide behavior were identified, providing valuable insights for future prevention and risk management planning in landslide-prone areas.

Keywords: Early warning index, Landslide, Physical modeling, Rainfall, Slope Stability of Soil

1) บทนำ

ปัญหาดินถล่มเป็นภัยพิบัติที่พบเห็นได้โดยทั่วไป ส่งผลกระทบต่อโครงสร้างพื้นฐาน เศรษฐกิจ และสิ่งแวดล้อม [1] ในพื้นที่ที่มีภูมิประเทศลาดชันซึ่งเกิดจากหลายปัจจัยได้แก่ การตัดไม้ทำลายป่า การเปลี่ยนแปลงสภาพภูมิอากาศ หรือการพัฒนาโครงสร้างพื้นฐานที่มีภูมิประเทศไม่เหมาะสม โดยเฉพาะในพื้นที่ที่มีเนินเขาหรือภูเขาที่ลาดชัน รวมถึงฝนตกหนักที่เกิดขึ้นบ่อยในบางพื้นที่ [2] เช่น ในปี ค.ศ. 2014 มีเหตุการณ์ดินถล่มในรัฐอูซิงตันที่เกิดจากฝนตกหนักและโครงสร้างดินที่ไม่เสถียรทำให้มีผู้เสียชีวิต 43 ราย [3] ในปี ค.ศ. 2017 เมืองโมโกอา ประเทศโคลอมเบีย เกิดเหตุดินถล่มจากฝนตกหนัก มีผู้เสียชีวิตมากกว่า 300 ราย [4] และจังหวัดชาวตะวันตก ประเทศอินโดนีเซียทำให้มีผู้เสียชีวิตหลายสิบคนเมื่อปี ค.ศ. 2021 [5] ซึ่งในบริเวณภาคเหนือตอนบนของประเทศไทย มักพบกับเหตุการณ์ดินถล่มบ่อยครั้งบริเวณเส้นทางคมนาคมตัดผ่านระหว่างอำเภอ เนื่องจากภูมิประเทศที่มีภูเขาสูงชันและมีการพัฒนาโครงสร้างพื้นฐานเพื่อเชื่อมระหว่างอำเภอ โดยตัดถนนบนภูเขา ภายได้เงื่อนไขที่จำกัดทางวิศวกรรม [6] ทำให้เมื่อมีฝนตกปริมาณมากจึงเกิดดินถล่มปิดทับเส้นทางเกิดความเสียหายอย่างมากต่อชีวิตและทรัพย์สิน การจราจร และสิ่งมีชีวิต [7] โดยเฉพาะเหตุการณ์ดินถล่มในบริเวณเส้นทางหลักในการเดินทางจากอำเภอแม่ริม เข้าสู่อำเภอสะเมิง จังหวัดเชียงใหม่ มักเกิดดินถล่มซ้ำ ๆ ทุกปี เช่น วันที่ 24 กันยายน ค.ศ. 2024 เกิดดินถล่มปิดเส้นทางสายกองขาหลวง - สะเมิงได้บริเวณดอยยาว [8] วันที่ 13 และ 25 กันยายน ค.ศ. 2022 ในพื้นที่ตำบลสะเมิงได้มีฝนตกสะสมทำให้เกิดดินถล่มปิดทางหลวงหมายเลข 1096 ช่วง กม. ที่ 24-25 ดังรูปที่ 1 [9], [10] ซ้ำกันในระยะเวลา 2 สัปดาห์ เจ้าหน้าที่ต้องใช้เครื่องจักรในการเปิดเส้นทางและตรวจสอบความปลอดภัย จะเห็นว่าดินถล่มมักเกิดขึ้นในช่วงฤดูฝนซึ่งมีปัจจัยเร่งสำคัญ คือปริมาณน้ำฝนที่ตกหนักเป็นระยะเวลานาน จนกระทั่งมวลดินอิ่มน้ำจนกระทั่งมีฝนตกหนักจนดินไม่สามารถสามารถอุ้มน้ำไว้ได้

ปัจจุบันหน่วยงานรัฐท้องถิ่นจะการใช้การเก็บสถิติความสัมพันธ์เหตุการณ์ดินถล่มและปริมาณน้ำฝนเฉลี่ยรายวันในรอบ 30 ปี ตามค่ามาตรฐานภูมิอากาศในการแจ้งเตือนพื้นที่ที่จะเกิดดินถล่มในระดับอำเภอ ซึ่งบางครั้งในพื้นที่การเกิดดินถล่มยังมีปัจจัยอื่น ๆ ร่วมด้วย เช่น ความลาดชัน การระบายน้ำ คุณสมบัติของดิน การรบกวนเสถียรภาพจากการตัดถนน เป็นต้น [11]

จากปัญหาดินถล่มทำให้ต้องมีการจัดการอย่างเหมาะสม การใช้เทคโนโลยีและมาตรการเชิงป้องกันมีบทบาทสำคัญในการลดความเสียหายที่อาจเกิดขึ้น โดยเฉพาะพื้นที่ที่มีความลาดชันสูงและมีโครงสร้างธรณีวิทยาที่แตกต่างกัน แนวทางการประเมินความเสี่ยงสามารถทำได้หลายวิธี เช่น การใช้แบบจำลองเชิงระบบสารสนเทศภูมิศาสตร์ (GIS-based models) และแบบจำลองเชิงตัวเลขเพื่อวิเคราะห์ความเสี่ยง [12] การใช้แบบจำลอง 3 มิติ (3D Seepage Modeling) ในการศึกษาปัจจัยที่มีผลต่อการเสถียรภาพของลาดดินในบริเวณที่มีสภาพไม่เอื้ออำนวย วิเคราะห์การเปลี่ยนแปลงของแรงดันน้ำในดินและผลกระทบต่อค่าความปลอดภัย (Factor of Safety: FoS.) อันเนื่องมาจากฝนตก [13] หรือระบบเตือนภัยโดยใช้เซ็นเซอร์ตรวจจับการเคลื่อนตัวของดิน [14] ซึ่งมีงานวิจัยที่ศึกษาเกี่ยวกับเกณฑ์ปริมาณฝนที่ส่งผลต่อเสถียรภาพของลาดดินเพื่อพัฒนาเกณฑ์เตือนภัยดินถล่มทั้งการใช้แบบจำลองทางกายภาพในการสังเกตและวิเคราะห์พฤติกรรมของดินในสถานการณ์ต่าง ๆ ได้อย่างเป็นระบบ [15]

การศึกษามูลค่าของปริมาณน้ำฝนในการเกิดดินถล่มด้วยแบบจำลองทางกายภาพ ช่วยให้เข้าใจพฤติกรรมของมวลดินและพัฒนาระบบเตือนภัยในพื้นที่มีถนนตัดผ่านบนภูเขาลาดชันที่ไม่สามารถออกแบบให้เหมาะสมต่อเสถียรภาพได้ ซึ่งกรณีศึกษาอยู่ที่เส้นทางหลวงหมายเลข 1096 เชียงใหม่-สะเมิง จะช่วยเติมเต็มองค์ความรู้ในบริบทของแนวทางพัฒนาการป้องกันและระบบเตือนภัยล่วงหน้าเพื่อลดความเสียหายในอนาคต



รูปที่ 1 : เหตุการณ์เกิดดินถล่มบริเวณเส้นทางหลวงหมายเลข 1096
เชียงใหม่ - สะเมิง [9]

2) ทบทวนวรรณกรรม

2.1) การศึกษาแบบจำลองพยากรณ์ดินถล่ม

การศึกษาแบบจำลองดินถล่มทางกายภาพในห้องปฏิบัติการเพื่อวิเคราะห์ปัจจัยปริมาณน้ำฝนและพฤติกรรมการเกิดดินถล่มเพื่อสร้างดัชนีการเตือนภัย ให้เฝ้าระวังความเสี่ยง สำหรับการเตรียมความพร้อมต่อเหตุการณ์ดินถล่มเพื่อลดความเสียหายที่จะเกิดขึ้นต่อชีวิตและทรัพย์สิน รวมถึงหาแนวทางการป้องกัน โดยส่วนใหญ่มักจะใช้การจำลองด้วยคอมพิวเตอร์วิเคราะห์หาผลกระทบต่อค่าความปลอดภัย [13] โดยใช้หลักฟิสิกส์ในการจำลองตัวเลขที่พิจารณาจากค่าพารามิเตอร์ด้านธรณีเทคนิค เช่น ความแข็งแรงเฉือน ความชื้น แรงดันน้ำ แบบจำลองส่วนใหญ่จะใช้วิธีสมดุลจำกัด (Limit Equilibrium Method: LEM) และวิธีไฟไนต์เอลิเมนต์ (Finite Element Method: FEM) ซึ่งช่วยให้สามารถคำนวณค่าปัจจัยความปลอดภัย และระบุตำแหน่งแนวการถล่มได้ [16] ทำให้สามารถวิเคราะห์สาเหตุการถล่มได้อย่างละเอียด หากแต่มีข้อจำกัดทางด้านความแม่นยำของข้อมูลทางธรณีเทคนิคและใช้ทรัพยากรในการประมวลผลสูง หรือการใช้แบบจำลอง GIS สามารถนำมาสร้างแผนที่ความไวต่อการเกิดดินถล่ม โดยผสมผสานข้อมูลเชิงพื้นที่ เช่น ความชื้น ธรณีวิทยา ปริมาณน้ำฝน และการใช้ที่ดิน แบบจำลอง GIS มีหลายวิธี เช่น แบบจำลองเชิงผู้เชี่ยวชาญ (heuristic models) แบบจำลองเชิงสถิติ (statistical models) และแบบจำลองการเรียนรู้ของเครื่อง (machine learning models) [17] ซึ่งสามารถวิเคราะห์พื้นที่ขนาดใหญ่ได้ ซึ่งเหมาะกับการประเมินในพื้นที่ขนาดใหญ่และยากต่อการเข้าถึง หากยังไม่สามารถใช้ในการพยากรณ์เวลาหรือขนาดของดินถล่มได้โดยตรง ทั้งนี้ยังขึ้นอยู่กับคุณภาพของข้อมูลและความแม่นยำของแบบจำลองที่ใช้

แบบจำลองทางกายภาพที่ใช้ในภาคสนามรวมถึงการตรวจวัดพื้นที่เสี่ยงดินถล่ม เช่น การติดตั้งระบบเตือนภัยโดยใช้เซ็นเซอร์ตรวจจับการเคลื่อนตัวของดิน เพื่อติดตามพฤติกรรมและปัจจัยแวดล้อมภาคสนาม [14] หรือการการกระตุ้นดินถล่มโดยควบคุมปริมาณน้ำฝน [18] สามารถควบคุมตัวแปรที่ใช้ในการศึกษาได้ดี ทำให้เข้าใจกลไกของดินถล่มได้ละเอียด แต่การสร้างอุปกรณ์และการดำเนินการมีค่าใช้จ่ายสูง รวมถึงความปลอดภัยในการดำเนินงานที่บริเวณพื้นที่ศึกษาค่อนข้างใช้เวลาในการศึกษาและเก็บข้อมูล การทดลองซ้ำทำได้ยากเนื่องจากได้รับผลกระทบจากตัวแปรที่เป็นสิ่งแวดล้อมรวมถึงมีพื้นที่ขนาดใหญ่อีกด้วย

ในการจำลองในห้องปฏิบัติการนั้นจะสามารถควบคุมตัวแปรได้ดี โดยเฉพาะสภาวะแวดล้อมที่คงที่ที่สามารถศึกษาสาเหตุและผลกระทบของดินถล่มที่เกิดจากปริมาณน้ำฝน ในรูปแบบของการทดลองในห้องปฏิบัติการที่เขต Hongkou ของเมือง Dujiangyan ประเทศจีน ได้จำลองจากพื้นที่ที่เกิดดินถล่มที่มีความชัน 28–40 องศา มีความกว้าง 416 เมตร ความยาว 190–260 เมตร ความลึก 4–10 เมตร ใช้มาตราส่วน 1:200 สำหรับความกว้างและความยาว และ 1:50 สำหรับความลึก ใช้ระบบจำลองฝนที่สามารถปรับความเข้มข้นได้และติดตั้งเซนเซอร์เพื่อวัดค่าทางกายภาพ ผลจากการจำลองการถล่มของดินพบว่าในช่วงปริมาณน้ำฝน 30–60 มิลลิเมตรต่อชั่วโมง แบบจำลองเกิดดินถล่มที่บริเวณส่วนปลายของดินลาดชัน และเมื่อมีฝนตกต่อเนื่องจะไม่เกิดการพังทลาย กระบวนการนี้เรียกว่า “Slow Retrogressive Toe Sliding” ช่วงที่ปริมาณน้ำฝน 90–150 มิลลิเมตรต่อชั่วโมง ส่วนที่ทรุดตัวลงที่ด้านล่างดินลาดชันจะค่อย ๆ กระจายตัวไปถึงตรงกลาง จนทำให้เกิดความสมดุลใหม่ กระบวนการนี้เรียกว่า “Multiple Retrogressive Toe Sliding” ช่วงที่ปริมาณน้ำฝน 180 มิลลิเมตรต่อชั่วโมง ขึ้นไปแบบจำลองลาดชันส่วนใหญ่จะเกิดการถล่มอย่างรวดเร็ว และส่วนที่เหลืออยู่จะถูกชะล้างโดยฝนที่ตกหนัก กระบวนการนี้เรียกว่า “Suddenly Diffusive Flow Sliding As a Whole” [15]

งานวิจัยนี้จึงมีแนวคิดในใช้แบบจำลองทางกายภาพเช่นเดียวกันเพื่อศึกษาพฤติกรรมของดินภายใต้ปริมาณน้ำฝนความเข้มข้น 4 ระดับที่ได้จากการเก็บสถิติในพื้นที่ย้อนหลัง 5 ปีและประเมินความเสี่ยงของดินถล่มบริเวณดินเดิมข้างทางที่ถูกตัดถนนผ่านและดินที่ถูกบดอัดเพื่อทำไหล่ทาง เพื่อสามารถสังเกตการเคลื่อนตัวของมวลดินและการเปลี่ยนแปลงโครงสร้างดินวิเคราะห์และประเมินความเสี่ยงด้วยวิธีดัชนีเตือนภัยล่วงหน้าเพื่อพัฒนาวิธีการป้องกันความเสี่ยงบริเวณที่อาจเกิดการถล่มได้

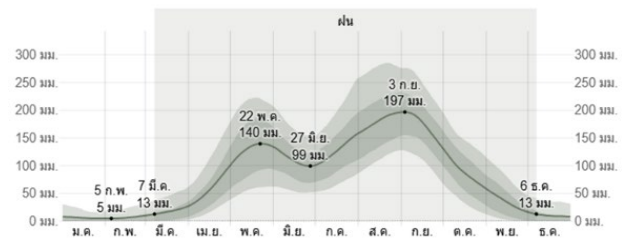
2.2) การศึกษาปริมาณน้ำฝนต่อการเกิดดินถล่ม

ในการศึกษาการเตือนภัยดินถล่มจากการตรวจวัดปริมาณน้ำฝนในพื้นที่ภูเขาทางภาคเหนือของประเทศไทย ได้ใช้การวิเคราะห์เชิงประจักษ์จากข้อมูลเหตุการณ์ดินถล่มในอดีตเพื่อหาความสัมพันธ์ระหว่างปริมาณฝนสะสม รายวันหรือรายชั่วโมงกับเหตุการณ์ดินถล่มพบว่า ฝนตกสะสมในช่วง 4 วัน มีความสัมพันธ์กับดินถล่มมากที่สุด โดย 3 วันแรกจะทำให้ดินอิ่มตัว และฝนวันที่ 4 เป็นปัจจัยกระตุ้นให้เกิดดินถล่ม นอกจากนี้ยังพบความแตกต่าง

ตามสภาพภูมิประเทศ จึงกำหนดเกณฑ์ฝนวิกฤตเฉพาะแต่ละโซน และพัฒนาเป็นแบบจำลองเตือนภัย เมื่อหากปริมาณฝนสะสมเกินเกณฑ์ก็ถือว่ามีความเสี่ยงสูงและควรแจ้งเตือนทันที [19] นอกจากนี้ยังมีการศึกษาใช้เกณฑ์ปริมาณฝนเชิงสถิติ เพื่อใช้เป็นแนวทางการเตือนภัยดินถล่มระดับประเทศ โดยพัฒนาจากข้อมูลฝนและเหตุการณ์ดินถล่มจำนวนมากทั่วประเทศนำมาวิเคราะห์ปริมาณฝนที่จะทำให้เกิดดินถล่มได้ [20] รวมถึงการใช้แผนที่พื้นที่เสี่ยงหรือความไวต่อดินถล่ม โดยเป็นวิธีการวิเคราะห์ข้อมูลเชิงพื้นที่แบบสองตัวแปร (bivariate) ในแบบจำลอง GIS คือการรวบรวมแผนที่ปัจจัยต่าง ๆ ในพื้นที่ศึกษาที่คาดว่าจะมีผลต่อดินถล่ม เช่น แผนที่ความชัน แผนที่ธรณีวิทยา แผนที่การใช้ที่ดิน เป็นต้น ประกอบกับแผนที่ตำแหน่งดินถล่มในอดีต จากนั้นจึงใช้การคำนวณความถี่อัตราการเกิดดินถล่มเพื่อสร้างระดับการเตือนภัยของดินถล่ม [21] อย่างไรก็ตามบริเวณเส้นทางหลวงหมายเลข 1096 เชียงใหม่-สะเมิง มักเกิดการดินถล่มในทุก ๆ ปี และมีการซ่อมแซมและป้องกันจุดที่เกิดเหตุดินถล่มและเฝ้าระวัง แต่ยังคงเกิดการถล่มบริเวณใหม่ตลอดเส้นทาง จึงเป็นที่น่าสนใจในการศึกษาเพิ่มเติมเกี่ยวกับพฤติกรรมมวลดินบริเวณดังกล่าวโดยการใช้แบบจำลองทางกายภาพเพื่อช่วยให้เข้าใจลักษณะพฤติกรรมของดินและการเคลื่อนตัวของดินลาดชันภายใต้ความเข้มข้นของฝนในระดับต่าง ๆ ที่อาจส่งผลกระทบต่อเปลี่ยนแปลงโครงสร้างของดินลาดชัน โดยพิจารณาทั้งปริมาณฝนและมวลดินลาดชัน 2 บริเวณที่มักจะเกิดเหตุการณ์ดินถล่มข้างทางและการหลุดตัวของถนนไหล่ทาง เพื่อเป็นแนวทางการสร้างระบบเตือนภัยในอนาคต

จากสถิติปริมาณน้ำฝนเฉลี่ย 5 ปี ย้อนหลัง ช่วงที่มีฝนชุกของปีมีระยะเวลา 9 เดือน ระหว่างเดือนมีนาคมถึงเดือนธันวาคม โดยช่วงที่สามารถวัดปริมาณน้ำฝนได้มากที่สุดของอำเภอสะเมิง คือเดือนสิงหาคม มีปริมาณน้ำฝนสูงสุดที่ 197 มิลลิเมตรต่อชั่วโมง ดังรูปที่ 2 ทำให้เห็นว่าช่วงเวลาที่ปริมาณน้ำฝนมากที่สุดมี 2 ช่วงเวลาของปี คือระหว่างเดือนมีนาคมถึงเดือนมิถุนายนเป็นช่วงเวลาของการสะสมตัวของน้ำในมวลดิน มีปริมาณน้ำฝนเฉลี่ย 100 มิลลิเมตรต่อชั่วโมง จนทำให้เกิดการอิ่มตัวของน้ำในดิน [12] และช่วงระหว่างเดือนมิถุนายนถึงเดือนกันยายนจะเป็นช่วงเวลาที่มักเกิดดินถล่ม เนื่องจากระดับความเข้มข้นน้ำฝนที่เพิ่มขึ้นในช่วงฤดูมรสุม [22] ปริมาณน้ำฝนสูงสุดในพื้นที่อยู่ที่ 270 มิลลิเมตรต่อชั่วโมงและต่ำสุดที่ 99 มิลลิเมตรต่อชั่วโมง [9] เป็นที่น่าสนใจอย่างยิ่งว่า เมื่อปริมาณน้ำฝนที่ระดับ 100, 160, 220 และ 270 มิลลิเมตรต่อชั่วโมง ลักษณะการเกิดดินถล่มจะทำให้สามารถ

วิเคราะห์และพัฒนาระบบเตือนภัยดินถล่มได้จากแบบจำลองภายใต้ปริมาณน้ำฝนได้



รูปที่ 2 : สถิติปริมาณน้ำฝนเฉลี่ยย้อนหลัง 5 ปี (พ.ศ. 2519–2023)
พื้นที่อำเภอสะเมิง จังหวัดเชียงใหม่ [19]

2.3) การวิเคราะห์เสถียรภาพและคุณสมบัติของดินลาดชัน

ลักษณะทางธรณีวิทยาของพื้นที่สะเมิงมีความซับซ้อน เนื่องจากมีการแทรกดันตัวของหินแกรนิตหลายครั้ง ทำให้หินเดิมถูกกระบวนการเปลี่ยนลักษณะและการแปรสภาพ หินที่พบในบริเวณนี้ประกอบด้วยหินแปรเกรดสูงที่มีอายุแตกต่างกัน ตั้งแต่ยุคแคมเบรียนถึงออร์โดวิเซียน ดินชั้นบน (Topsoil) มีลักษณะเฉพาะที่สัมพันธ์กับธรณีวิทยาของพื้นที่ จากการศึกษาลักษณะดินในป่าดิบเขาต่ำที่เหลือเป็นหย่อม ดินชั้นบนมีความหนาแน่นรวมต่ำและมีเนื้อหยาบเป็นดินร่วนปนทราย ส่วนดินชั้นล่างเป็นดินร่วนเหนียวปนทราย ดินร่วน และร่วนเหนียว อีกทั้งมีรอยเลื่อนมีพลังในพื้นที่สะเมิงส่งผลกระทบต่อโครงสร้างทางธรณีวิทยาและสภาพดินในพื้นที่ซึ่งอาจเพิ่มความเสี่ยงต่อการเกิดดินถล่มที่เกิดจากการเคลื่อนตัวของมวลดินและหิน โดยมีปัจจัยลักษณะทางธรณีวิทยาเป็นบริเวณที่มีหินผุให้ชั้นดินหนา สภาพภูมิประเทศเป็นพื้นที่ภูเขาสูงและมีความลาดชัน ลักษณะสิ่งแวดล้อมเปลี่ยนแปลงจากการใช้ประโยชน์ที่ดินโดยไม่ถูกหลักวิชาการ เช่นการตัดถนนผ่านภูเขาสูง และปริมาณน้ำฝนที่มากจนชั้นดินอุ้มน้ำไม่ไหว [23]

โดยส่วนใหญ่ดินมีลักษณะการวิบัติแบบรูปโค้ง (circular failure) ซึ่งสามารถวิเคราะห์ด้วยชุดของแผนภูมิ แผนภาพการวิบัติแบบโค้งของลาดดิน (circular failure chart) โดยใช้หาปริมาตรมวลดินถล่ม ใช้จุดศูนย์กลางบนส่วนโค้งของระนาบการวิบัติ (X, Y) ตำแหน่งที่เกิดรอยแตกแบบตึง บริเวณด้านบนของมวลดินลาดชันที่ได้จากแบบจำลองทางกายภาพ ซึ่งพิกัด (X, Y) จะวัดโดยมีจุดอ้างอิงอยู่ที่ส่วนล่างสุดของดินลาดชัน และระยะ b ที่เกิดรอยแตกแบบตึง นำมาประมาณการค่ามุมเสียดทานภายในของดิน (friction angle) [24] และคำนวณปริมาตรของมวลดิน

ที่เกิดถล่ม เพื่อใช้ในการประเมินความเสี่ยงด้วยวิธีดัชนีเตือนภัยล่วงหน้า

2.4) การประเมินความเสี่ยงด้วยวิธีการดัชนีเตือนภัยล่วงหน้า

ระดับความรุนแรงของเหตุการณ์ดินถล่มสามารถจัดจำแนกได้เป็นสองด้านหลัก ๆ ได้แก่ ระยะเวลาของฝนตก และขนาดของพื้นที่ที่เกิดการถล่ม ตัวอย่างเช่น การเกิดมวลดินถล่มที่มีขนาดใหญ่เนื่องจากจากฝนตกในระยะเวลาสั้น ๆ สามารถอธิบายได้ว่าเป็นเหตุการณ์ที่มีความรุนแรงค่อนข้างสูง นอกจากนี้ระดับความรุนแรงยังสามารถกำหนดได้โดยการเปรียบเทียบความเสียหายกับเหตุการณ์ที่เกิดขึ้นในสภาพธรณีวิทยาเดียวกันเนื่องจากปริมาณของฝนที่แตกต่างกัน [15]

ระดับความรุนแรงของความเสียหายจากดินถล่มที่เกิดจากปริมาณความเข้มข้นระดับน้ำฝนที่ตกและระยะเวลาของฝนที่ตกลงมา (กล่าวคือเวลาที่ให้เกิดมวลดินถล่ม) จะแทนด้วย de และ T (นาทีก) ตามลำดับ อัตราส่วนของปริมาตรของมวลดินที่ถล่มต่อปริมาตรของแบบจำลองดินลาดชันจะแทนด้วย V และ $|T|$ เป็นค่าที่ไม่มีมิติของ T ตามการร่อนมาข้างต้น ยิ่งระยะเวลาฝนตกสั้นและปริมาตรมวลที่ดินถล่มมากขึ้น ความเสียหายจะยิ่งรุนแรงมากขึ้น ดังนั้น de จึงเป็นสัดส่วนโดยตรงกับ V แต่เป็นสัดส่วนผกผันกับ $|T|$ ซึ่งสามารถแสดงได้ดังนี้

$$de = \frac{V}{|T|} \quad (1)$$

โดยที่ V คือ อัตราส่วนของปริมาตรที่พังทลาย

$|T|$ คือ เวลาที่เกิดการพังทลาย

ดังนั้นระดับความรุนแรงของเหตุการณ์ดินถล่มเนื่องจากปริมาณของฝนที่ I : มิลลิเมตรต่อชั่วโมง สามารถกำหนดได้ดังนี้

$$R = \frac{de_i}{de_{max}} \quad (2)$$

โดยที่ de_i คือ อัตราส่วนของปริมาตรที่มวลดินถล่มต่อเวลาที่เกิดดินถล่ม

de_{max} คือ อัตราส่วนของปริมาตรที่มวลดินถล่มต่อเวลาที่เกิดดินถล่มสูงสุด

ดังนั้น ค่า V , de และ R ที่คำนวณได้ นำไปวิเคราะห์อัตราส่วนดินถล่มเนื่องจากปริมาณของฝนแต่ละระดับ โดยแสดงในตารางที่ 1 [15]

ตารางที่ 1 : ระดับความรุนแรงของดินถล่มเนื่องจากปริมาณฝน

สัญลักษณ์	ปริมาณน้ำฝน (มิลลิเมตรต่อชั่วโมง)		
	i_0	i_1	i_n
V	V'_{i0}/V	V'_{i1}/V	V'_{in}/V
$ T $	$ T _0$	$ T _1$	$ T _n$
de	de_{i0}	de_{i1}	de_{max}
R	R_{i0}	R_{i1}	R_{in}

3) วิธีดำเนินการวิจัย

3.1) การเก็บข้อมูลพื้นที่เพื่อนำมาใช้ในการทดลอง

การเก็บข้อมูลจากพื้นที่แขวงทางหลวงเชียงใหม่ 2 ทล.1096 เชียงใหม่-สะเมิง กม. ที่ 24-25 โดยใช้หน่วยน้ำหนักของดินที่ได้จากการทดสอบภาคสนาม 2 บริเวณ คือ พื้นที่บริเวณดินเดิมข้างทางที่ถูกถนนตัดผ่าน โดยเลือกขึ้นไปเก็บบริเวณด้านบน และดินบดอัดข้างถนนบริเวณไหล่ทาง พื้นที่การศึกษามีลักษณะการออกแบบถนนที่พิจารณาพร้อมกับความลาดเอียงของภูมิประเทศที่มักใช้ในการตัดทางทำถนนบนพื้นที่เข้าสู่ชันทางภาคเหนือ จากการวัดความชันลาดมวลดินเท่ากับ 70 องศา และการออกแบบถนนด้วยวิธี Cut/Fill (พิจารณาร่วมกับสภาพความลาดเอียงของภูมิประเทศ) [25] ดังรูปที่ 3



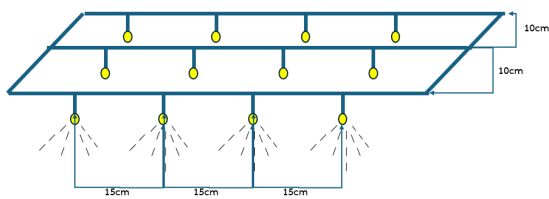
รูปที่ 3 : บริเวณที่เก็บตัวอย่างดิน และเก็บข้อมูลภาคสนาม

ช่วงเวลาเก็บข้อมูลคือเดือนพฤษภาคม ค.ศ. 2023 ซึ่งเป็นช่วงที่เข้าสู่ฤดูฝน ตัวอย่างดินที่เก็บได้จึงมีความชื้นสูง จากการทดสอบหน่วยน้ำหนักมวลดินเทียบกับ Proctor Test ในห้องปฏิบัติการ พบว่าดินเดิมข้างทางและดินบดอัดไหล่ทางมีค่าเท่ากับ 15.30 kN/m^3 และ 18.34 kN/m^3 ตามลำดับ มีค่าความชื้นของดิน (water Content) ค่าประมาณ 18% เท่ากับค่าความชื้นของ

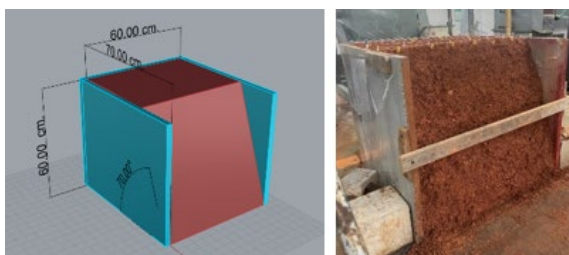
มวลดินที่อิ่มตัวด้วยน้ำ จากการทดสอบการบดอัด (compaction test) ในห้องปฏิบัติการ จึงใช้เป็นแบบจำลองมวลดินที่อิ่มน้ำ เพื่อศึกษาดินถล่มภายใต้ฝนตกที่ความเข้มข้นระดับ 4 ระดับ ได้แก่ 100, 160, 220 และ 270 มิลลิเมตรต่อชั่วโมง เนื่องจากเดือนธันวาคม ถึงเดือนมิถุนายน เป็นช่วงของการสะสมตัวของน้ำในดิน จนทำให้ดินเกิดการอิ่มตัว และการเกิดดินถล่มบริเวณทางหลวงหมายเลข 1096 มักอยู่ในช่วงฤดูฝนระหว่างเดือนมิถุนายน ถึงเดือนพฤศจิกายน โดยเมื่อ กันยายน ค.ศ. 2023 มีปริมาณน้ำฝนต่ำสุดที่ 99 มิลลิเมตรต่อชั่วโมง และสูงสุดที่ 270 มิลลิเมตรต่อชั่วโมง

3.2) แบบจำลองทางกายภาพและการควบคุมปริมาณน้ำฝน

พื้นที่อำเภอสะเมิง จังหวัดเชียงใหม่ มีธรณีวิทยาซับซ้อน
ชั้นดินอ่อนและสภาพฝนตกต่อเนื่อง การใช้แบบจำลองกายภาพ
เพื่อให้สามารถจำลองลักษณะภูมิประเทศ หน่วยน้ำหนักมวลดิน
ความชื้น และฝนตกได้ใกล้เคียงกับสภาพในพื้นที่จริง โดยต้องการ
เข้าใจพฤติกรรมของดินลาดชันและลักษณะการเคลื่อนตัวแบบ
real-time เช่น การแตกร้าว ลักษณะดินถล่ม หรือการไหลของ
มวลดิน ซึ่งสามารถเติมเต็มช่องว่างของแบบจำลองเชิงตัวเลขที่
อาศัยความแม่นยำในการป้อนข้อมูลและการทำนายตำแหน่งที่
เกิดมวลดินถล่ม ซึ่งจะช่วยสร้างระบบเตือนภัยที่เข้าใจง่ายและใช้
งานได้จริง โดยสร้างแบบจำลองขนาดสูง 60 เซนติเมตร ความ
กว้าง 60 เซนติเมตร ความลึก 70 เซนติเมตร มีความชันที่ 70
องศา ดังรูปที่ 4 และ 5



รูปที่ 4 : แบบจำลองทางกายภาพและระบบจำลองปริมาณน้ำฝน

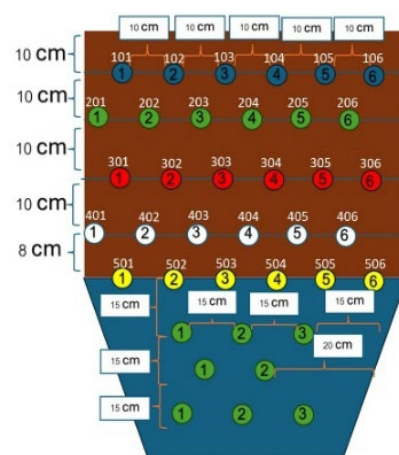


รูปที่ 5 : แบบจำลองทางกายภาพและระบบจำลองปริมาณน้ำฝน

ปริมาตรของแบบจำลองเท่ากับ 0.2127 ลูกบาศก์เมตร ใช้ตัวอย่างดินที่เก็บจากพื้นที่ กำหนดค่าความชื้นประมาณ 18% ใช้แบบจำลองมวลดินที่มีหน่วยน้ำหนัก 15.30 kN/m^3 และ 18.34 kN/m^3 สำหรับสถานการณ์ที่ฝนตก 4 ระดับ ระบบน้ำฝนใช้หัวพ่นหมอกเพื่อให้น้ำฝนกระจายตัวทั่วพื้นที่และสามารถควบคุมปริมาณความเข้มของน้ำฝนได้ เพื่อศึกษาพฤติกรรมการเกิดดินถล่มเนื่องจากฝนตกในปริมาณที่ดินไม่สามารถอุ้มน้ำได้

3.3) การวัดค่าการเคลื่อนตัวและลักษณะดินถล่ม

แบบจำลองจะถูกรวบรวมชุด 5 แถว บริเวณด้านบนบนและด้านหน้าของดินลาดชัน โดยวางห่างกัน 10 เซนติเมตร และวางแบบสลับฟันปลาในแต่ละแถว เพื่อใช้เป็นตำแหน่งอ้างอิงในการเก็บข้อมูลการเคลื่อนตัวของหมุดทั้งในแนวระนาบ และแนวตั้ง ดังแสดงในรูปที่ 6 การเก็บข้อมูลการเคลื่อนตัวของมวลดินแต่ละจุดคือตั้งแต่ช่วงขณะที่เกิดฝนตก และช่วงที่หลังฝนหยุดตก โดยเก็บข้อมูลการเคลื่อนตัวทุก ๆ 30 นาที และช่วงระยะ 3 ชั่วโมง จนครบระยะเวลา 13 ชั่วโมงนับตั้งแต่เริ่มปล่อยน้ำฝนพร้อมทั้งตั้งกล้องบันทึกภาพวิดีโอ ซึ่งในการศึกษาการเกิดดินถล่มในไต้หวัน จะใช้สถิติ Descriptive และ Time Series Analysis พบว่า 80% ของดินถล่มทั้งหมดเกิดภายใน 10 ชั่วโมงหลังฝนหยุดตก โดยเฉพาะในพื้นที่ลาดชันสูงที่มีฝนตกหนักมาก่อนหน้า [26] แสดงถึงพฤติกรรมมวลดินมีการเคลื่อนที่อย่างต่อเนื่องหลังฝนตก ซึ่งอาจเป็นพฤติกรรมการคืบของดิน



รูปที่ 6 : การกำหนดตำแหน่งอ้างอิงเพื่อเก็บข้อมูลการเคลื่อนตัว

การวิเคราะห์ปริมาณดินถล่มของแบบจำลอง จะเป็นรูปแบบ
การพังทลายรูปโค้ง ในการคำนวณค่าระยะ X , Y จุดศูนย์กลาง
ของส่วนโค้ง จากระยะรอยแตกแบบดึง (Tension Crack: b) ทำ

การวัดระยะรอยแยกดินที่ปรากฏบนแบบจำลองดินลาดชัน โดยใช้จุดศูนย์กลางของส่วนโค้งที่เกิดการวิบัติ [27] และตำแหน่งการยุบตัวในแนวตั้งของพื้นผิวด้านบนของแบบจำลอง นำไปขึ้นรูป 3 มิติ เพื่อใช้หาปริมาตรดินที่เกิดการถล่มจากน้ำฝนระดับต่าง ๆ

3.4) การประมาณการอัตราการถล่มของมวลดิน

การคำนวณปริมาตรดินถล่ม นำค่าปริมาตรที่เกิดดินถล่มมาเปรียบเทียบกับปริมาตรเดิมของแบบจำลองมวลดินลาดชัน จะได้ค่า V ซึ่งหมายถึงอัตราการพังทลาย ตามสมการ 3

$$V = \frac{V_i}{V_0} \quad (3)$$

โดยที่ V_i คือ ปริมาตรที่เกิดดินถล่ม (m^3)

V_0 คือ ปริมาตรเดิมก่อนเกิดดินถล่ม (m^3)

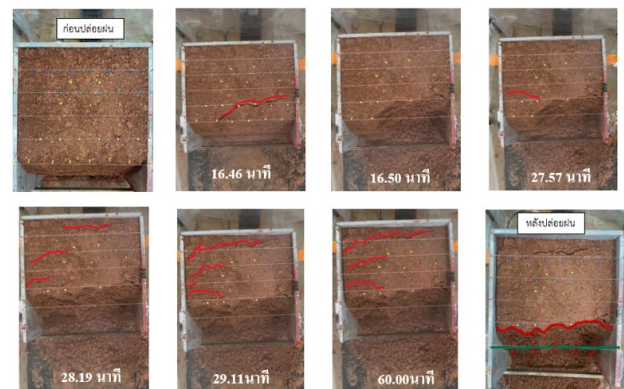
คำนวณค่า d_e ความเข้มของปริมาณน้ำฝนที่มีผลต่อดินถล่ม (effective rainfall depth) และคำนวณค่า R ซึ่งหมายถึงค่าความรุนแรงของอัตราดินถล่มที่เกิดจากปริมาณน้ำฝน หลังจากนั้นนำมาพัฒนาตารางดัชนีเตือนภัยล่วงหน้าดินถล่มที่เกิดจากปริมาณน้ำฝนและระดับความรุนแรงของดินถล่มที่เกิดขึ้น [15]

4) ผลการวิจัยและอภิปรายผล

4.1) พฤติกรรมของดินถล่มเมื่อเกิดฝนตก

จากการบันทึกพฤติกรรมแบบจำลองมวลดินลาดชัน โดยติดตั้งกล้องบันทึกวิดีโอและวัดระยะการเคลื่อนตัวของหมุดอ้างอิง เมื่อปล่อยน้ำฝนที่ระดับต่างกันพบว่า ลักษณะการเกิดดินถล่มของแบบจำลองดินลาดชันทั้ง 2 ชนิด พบว่า ดินเดิมบริเวณข้างทางที่มีหน่วยน้ำหนัก 15.30 kN/m^3 เมื่อมีฝนตกลงมาปริมาณตั้งแต่ 100 มิลลิเมตรต่อชั่วโมง สังเกตได้ว่ามีน้ำซึมออกบริเวณด้านหน้าของมวลดินลาดชัน และเกิดรอยแยกดินบริเวณด้านบนเพิ่มขึ้นอย่างช้า ๆ และใช้เวลามากกว่า 10 นาที ดังรูปที่ 7 จากนั้นจึงเกิดดินถล่มบริเวณส่วนขอบด้านบนของแบบจำลองก่อน และเมื่อมีฝนตกต่อเนื่องด้านล่างจะไม่เกิดการถล่ม กระบวนการนี้เรียกว่า “Slow Retrogressive Crest Sliding” เมื่อทดสอบในแบบจำลองที่ปล่อยระดับปริมาณน้ำฝนมากขึ้นพบว่า ปรากฏส่วนที่ดินถล่มเป็นบริเวณที่ขอบด้านบนถึงส่วนล่างของแบบจำลอง โดยมวลดินลาดชันจะค่อย ๆ ทรุดกระจายตัวอย่างต่อเนื่องไปถึงส่วนกลาง จนทำให้เกิดความสมดุลใหม่ กระบวนการนี้เรียกว่า “Multiple Retrogressive Sliding”

สำหรับดินบดอัดไหลทางที่มีหน่วยน้ำหนัก 18.34 kN/m^3 จะเกิดดินถล่มที่ระดับปริมาณน้ำฝน 220 มิลลิเมตรต่อชั่วโมงขึ้นไป มีการสะสมของน้ำในดินเพิ่มขึ้นอย่างมาก แต่พบการซึมผ่านของน้ำด้านหน้าดินลาดชันน้อย แบบจำลองเกิดรอยแยกด้านบนและเกิดการถล่มอย่างรวดเร็วทันทีทันใด เมื่อเวลาผ่านไปประมาณ 44 นาที ส่วนที่เหลืออยู่จะถูกชะล้างโดยฝนที่ตกหนัก กระบวนการนี้เรียกว่า “Suddenly Diffusive Flow Sliding as a Whole” [28] ดังรูปที่ 8

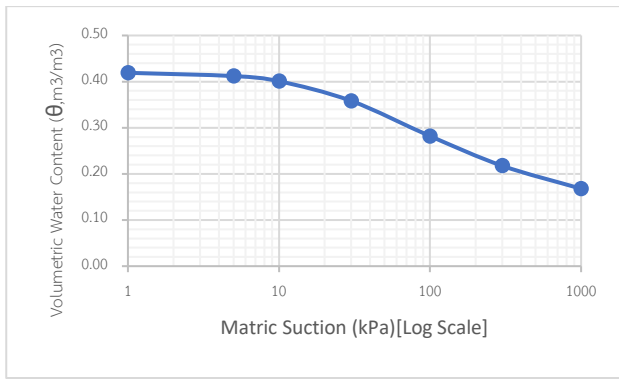


รูปที่ 7 : การพังทลายของแบบจำลองมวลดินเดิมที่หน่วยน้ำหนัก 15.30 kN/m^3 ที่ระดับน้ำฝน 100 มิลลิเมตรต่อชั่วโมง



รูปที่ 8 : การพังทลายของแบบจำลองมวลดินบดอัดที่หน่วยน้ำหนัก 18.34 kN/m^3 ระดับน้ำฝน 220 มิลลิเมตรต่อชั่วโมง

การวิเคราะห์การซึมผ่านของน้ำสำหรับแบบจำลองนี้โดยใช้สมการ Van Genuchten สำหรับดินชนิด ดินเหนียวปนทราย (sandy clay loam) ที่ได้จากการวิเคราะห์ขนาดเม็ดดิน และการทดสอบค่า Compaction Test ในห้องปฏิบัติการ แสดงอัตลักษณ์น้ำในดิน (soil water retention curve) สะท้อนพฤติกรรมการคายน้ำของดิน ดังรูปที่ 9



รูปที่ 9 : กราฟแสดงอัตลักษณ์น้ำในดิน (SWRC)

กราฟ SWRC ของดินเหนียวปนทรายมีลักษณะโค้งงอลง ชัดเจน (concave downward) บ่งบอกว่าดินอุ้มน้ำได้มากเมื่อแรงดูดน้ำ (matric suction) ยังต่ำ แต่เมื่อแรงดูดเพิ่มขึ้น ความชื้นโดยปริมาตรจะลดลงอย่างรวดเร็ว ซึ่งสอดคล้องกับพฤติกรรมดินไม่อิ่มตัวทั่วไป ในขั้นแรก น้ำจะถูกกักอยู่ในรูพรุนขนาดใหญ่จนกระทั่งเกิดแรงดูดเล็ก ๆ (air entry value) แล้วจึงเริ่มปล่อยน้ำออก เมื่อแรงดูดผ่านช่วง 10–100 kPa ดินจะปล่อยน้ำได้เร็วที่สุดสะท้อนถึงการระบายออกของรูพรุนขนาดกลางเป็นหลัก ก่อนที่แรงดูดจะสูงขึ้นจนเหลือเพียงน้ำตกค้างในรูพรุนจิ๋ว θ_r Residual Water Content ซึ่งแทบไม่เปลี่ยนแปลงอีกต่อไป

ในเชิงกลศาสตร์ของดิน ค่าความชื้นอิ่มตัว (θ_s) ประมาณ $0.42 \text{ m}^3/\text{m}^3$ บ่งชี้ว่าดินยังสามารถรับน้ำได้มากหลังฝนตกแทรกเข้าสู่พื้นผิว และค่าความชื้นตกค้าง (θ_r) ประมาณ $0.08 \text{ m}^3/\text{m}^3$ บ่งชี้ว่าดินยังเก็บน้ำได้แม้แรงดูดสูง การรู้ช่วง Field Capacity ($\theta \approx 0.35\text{--}0.40$) ช่วยให้คาดการณ์จังหวะที่แรงดันน้ำในรูพรุน (u) จะพุ่งขึ้น ลดแรงเฉือนสุทธิและนำไปสู่การเกิดดินถล่ม ส่วนช่วง Drainage Zone (10–100 kPa) เป็นช่วงวิกฤตที่หากดินถูกเพิ่มน้ำเข้าสู่บริเวณนี้ก็จะเกิดการเปลี่ยนแปลงอย่างอย่างรวดเร็ว

ดังนั้น กราฟ SWRC ช่วยอธิบายแบบจำลองฝนตกว่า ดินที่มีหน่วยน้ำหนักเบาจะถล่มเมื่อมีปริมาณฝนตกเข้มข้นเพิ่มเป็น 220–270 มิลลิเมตรต่อชั่วโมง ดินจะลด Matric Suction ลงสู่ช่วงนี้ภายใน 8–10 นาที เมื่อแรงดูดลดลงเข้าสู่ช่วง Transition Zone ขณะที่ดินที่มีหน่วยน้ำหนักมากต้องใช้เวลานานกว่าในการสะสมแรงดันรูพรุนให้ถึงระดับวิกฤตก่อนจะสั่นไหวออกไป แม้ปริมาณน้ำฝนสูงถึง 220 มิลลิเมตรต่อชั่วโมง จะใช้เวลาถึง 44 นาทีจึงถล่ม เพราะต้องรอให้ Matric Suction ลดผ่านช่วง Transition Zone จึงเข้าสู่สภาวะดินถล่มจริง

จากการเก็บบันทึกและวิเคราะห์ลักษณะดินถล่มของแบบจำลองมวลดิน โดยที่ดินเดิมข้างทางมีหน่วยน้ำหนัก 15.30 kN/m^3 และดินบดอัดไหล่ทางที่มีหน่วยน้ำหนัก 18.34 kN/m^3 ทดสอบภายใต้ปริมาณน้ำฝนแตกต่างกัน ทั้งหมด 8 แบบจำลอง สามารถจำแนกลักษณะพฤติกรรมการถล่มของดินลาดชัน ออกเป็น 4 ลักษณะ ดังนี้

4.1.1) เกิดการถล่มบริเวณด้านหน้าดินลาดชันที่ชะงักอย่างช้า ๆ ผลกระทบของฝนที่ตกต่อเนื่องไม่ได้ทำให้เกิดดินถล่มเพิ่มเติมนวลดินบางส่วนด้านบนไหลลงไปยังด้านล่างช้าๆ เรียกว่า “การเกิดดินถล่มอย่างช้าๆ บางส่วนของดินลาดชัน”

4.1.2) เกิดการถล่มบริเวณส่วนกลางของดินลาดชัน มีการไหลในแต่ละส่วนด้านบน เกิดขึ้นในช่วงเวลาที่ห่างกัน ลักษณะพฤติกรรมดินถล่มนี้เรียกว่า “การเกิดดินถล่มหลายครั้งอย่างต่อเนื่อง”

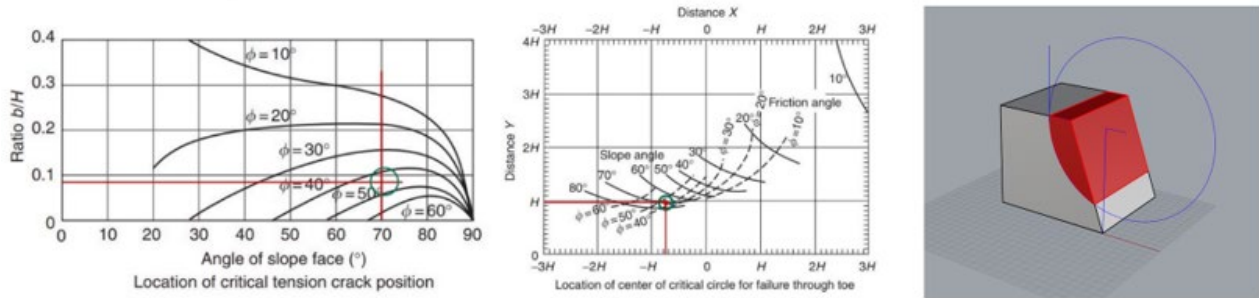
4.1.3) เกิดการถล่มบริเวณส่วนกลางของดินลาดชันและต่อเนื่อง เกิดการไหลตั้งแต่หมดแถวที่ 3 การถล่มเกิดขึ้นในแต่ละส่วนในช่วงเวลาที่ห่างกันและต่อเนื่อง ลักษณะพฤติกรรมนี้เรียกว่า “การเกิดดินถล่มหลายครั้งอย่างต่อเนื่อง”

4.1.4) เกิดการถล่มบริเวณส่วนกลางของดินลาดชันและเกิดขึ้นอย่างรวดเร็ว เนื่องจากการสะสมตัวของแรงภายใน ทำให้มีการเลื่อนถล่มของมวลดินอย่างรวดเร็วและมีการไหลอย่างต่อเนื่องในช่วงเวลาที่ยังมีฝนตกระยะเวลานึง

ลักษณะการเคลื่อนตัวของมวลดินที่ได้จากแบบจำลองพบว่า ส่วนใหญ่จะเกิดขึ้นบริเวณส่วนบนของดินลาดชัน ซึ่งมีลักษณะใกล้เคียงกับพื้นที่เกิดเหตุการณ์ดินถล่มเส้นทางหลวง 1096 อำเภอสะเมิง จังหวัดเชียงใหม่ โดยดินเดิมบริเวณข้างทางที่มักจะไหลลงมาปิดทับเส้นทางและดินบดอัดไหล่ทางมักจะทรุดลงเหวรวมถึงบางส่วนของถนนอีกด้วย

4.2) การวิเคราะห์พฤติกรรมดินถล่มเชิงสถิติและการประเมินความเสี่ยง

ปริมาณการดินที่เกิดดินถล่ม ได้จากการบันทึกค่าการเปลี่ยนแปลงของแต่ละหมุดในแนวระดับ นำไปสร้างแบบจำลองหลังการทดสอบด้วยโปรแกรมคอมพิวเตอร์ แสดงในรูปที่ 10 โดยจำลองการถล่มของมวลดินรูปแบบพังทลายรูปโค้ง จากระยะรอยแตกแบบตึง ใช้ชุดแผนภูมิประกอบเพื่อหาค่ามุมเสียดทานภายในและระยะจุดศูนย์กลางของส่วนโค้งที่เกิดดินถล่ม X และ Y ใช้ในการคำนวณประมาณการปริมาตรของดินที่ถล่มหลังฝนตก



รูปที่ 10 : การหาค่าระยะ X และระยะ Y และแบบจำลองทางคอมพิวเตอร์เพื่อหาปริมาตรดินถล่ม

จากการศึกษาพบว่า บริเวณด้านบนของแบบจำลอง หมดแถวที่ 3-5 เมื่อมวลดินเริ่มได้รับน้ำฝนในแต่ละระดับที่ปริมาณความเข้มข้นสูงขึ้น จะเกิดรอยแยกถึงบริเวณด้านบนและลึกเข้ามาในมวลดินมากขึ้น ทำให้ปริมาตรมวลดินที่เกิดถล่มเพิ่มมากขึ้นตามระดับน้ำฝนที่เพิ่มขึ้น ก่อนมวลดินถล่มตัวอย่างช้า ๆ ซึ่งสามารถนำการบันทึกค่าการเคลื่อนตัวในช่วงเวลาต่าง ๆ เพื่อหาอัตราการเคลื่อนตัวของดินลาดชันตั้งแต่เกิดรอยแยกถึงจนกระทั่งเกิดดินถล่ม ซึ่งสามารถนำไปศึกษาพฤติกรรมการคืบตัวของมวลดินได้อีกด้วย

ในการทดลองจะมีความแตกต่างของระดับปริมาณน้ำฝน คือ 100 160 220 และ 270 มิลลิเมตรต่อชั่วโมง ต่อแบบจำลองมวลดินทั้ง 2 แบบ เพื่อสังเกตพฤติกรรมและเวลาที่เกิดดินถล่มที่แสดงในตารางที่ 2

ตารางที่ 2 : ระยะเวลาการพังทลายของแบบจำลอง

ปริมาณน้ำฝน (มิลลิเมตรต่อ ชั่วโมง)	เวลาที่เกิดดินถล่ม (นาท)	
	มวลดินเดิม หน่วย น้ำหนัก 15.30 kN/m ³	มวลดินบดอัดใหม่ ทาง 18.34 kN/m ³
100	16.50	ไม่เกิดการถล่ม
160	14.03	ไม่เกิดการถล่ม
220	10.04	45.27
270	8.16	42.34

แบบจำลองมวลดินทั้ง 2 หน่วยน้ำหนัก มีความสัมพันธ์แบบยกกำลังสอง (power law) ของปริมาณระดับน้ำฝนกับระยะเวลาที่เกิดดินถล่ม ดังรูปที่ 11 แสดงให้เห็นว่า D ลดลงเมื่อปริมาณฝนเพิ่มขึ้น โดยที่ดินเดิมข้างทางที่มีหน่วยน้ำหนัก 15.30 kN/m³ มีการลดลงของเวลาที่เกิดดินถล่มมากกว่าที่ดินที่มีหน่วยน้ำหนัก 18.34 kN/m³ ค่าแนวโน้มระยะเวลาที่将会เกิดดินถล่มมีความ

สัมพันธ์แบบผกผัน (Inverse Relationship) ระหว่าง D และ I ดังสมการที่ 4 และ 5 ตามลำดับ

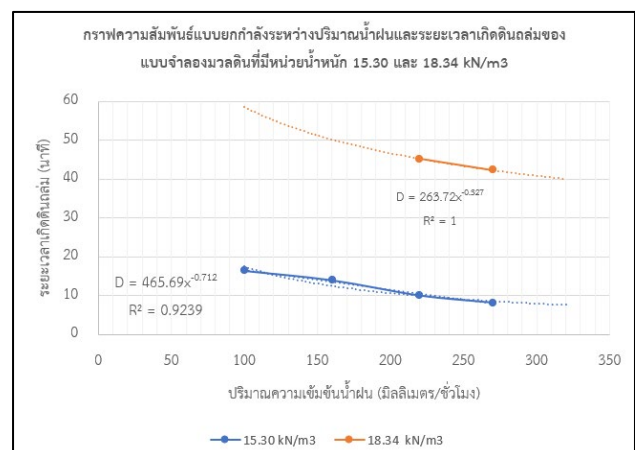
$$D = 465.69I^{-0.712} \quad (4)$$

$$D = 263.72I^{-0.327} \quad (5)$$

โดยที่ D คือ ระยะเวลา

I คือ ปริมาณน้ำฝน

จากความสัมพันธ์ดังกล่าวแบบจำลองสามารถประมาณการระยะเวลาการเกิดดินถล่มภายใต้ปริมาณฝนความที่มีความเข้มข้นระดับต่างๆ ได้ ซึ่งค่าสัมประสิทธิ์ความสัมพันธ์กำลังสอง (R^2) เข้าใกล้ 1 สะท้อนว่าแบบจำลองมีความสอดคล้องกับข้อมูลทดลอง และสามารถนำไปใช้ได้กับทั้งมวลดินข้างทางที่มีหน่วยน้ำหนัก 15.30 kN/m³ และดินบดอัดที่มีหน่วยน้ำหนัก 18.34 kN/m³ ในสถานการณ์ที่มวลดินอิ่มน้ำในฤดูฝน



รูปที่ 11 : ความสัมพันธ์แบบยกกำลังระหว่างปริมาณน้ำฝน (มิลลิเมตรต่อชั่วโมง) และระยะเวลาเกิดดินถล่ม (D นาท)

4.3) การประเมินความเสี่ยงด้วยวิธีดัชนีเตือนภัยล่วงหน้า

ในการประเมินดินถล่มของแบบจำลอง ค่า V (อัตราดินถล่ม) $|T|$ (เวลาที่เกิดดินถล่ม) de (ปริมาณน้ำฝนที่มีผลต่อดินถล่ม) และ R (ความรุนแรงของอัตราดินถล่มที่เกิดจากปริมาณน้ำฝน) ซึ่งใช้สมการที่ (1-4) ในการคำนวณ แสดงในตารางที่ 4 และ 5

ตารางที่ 4 : การประเมินดินถล่มที่หน่วยน้ำหนัก 15.30 kN/m³

สัญลักษณ์	ปริมาณน้ำฝน (มิลลิเมตรต่อชั่วโมง)			
	100	160	220	270
V	0.07	0.0834	0.132	0.1408
$ T $	16.50	14.03	10.04	8.16
De	0.0042	0.0059	0.013	0.017
R	0.246	0.345	0.762	1

ตารางที่ 5 : การประเมินดินถล่มที่หน่วยน้ำหนัก 18.34 kN/m³

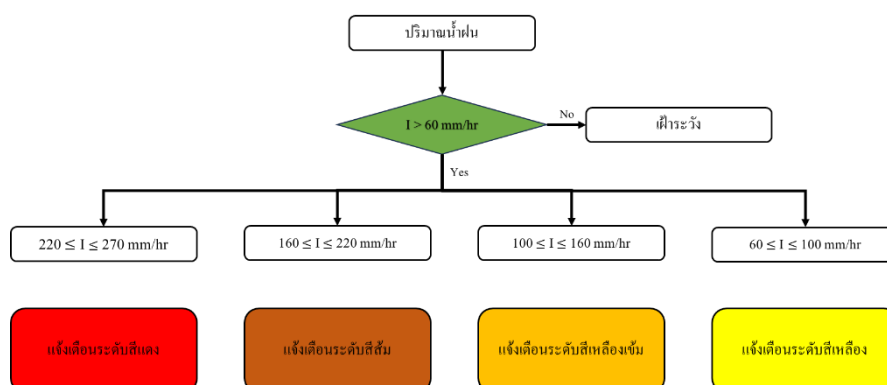
สัญลักษณ์	ปริมาณน้ำฝน (มิลลิเมตรต่อชั่วโมง)			
	100	160	220	270
V	-	-	0.0834	0.116
$ T $	-	-	45.27	42.34
De	-	-	0.0020	0.0027
R	-	-	0.72	1

จากตารางแสดงให้เห็นว่าปริมาณน้ำฝนที่สูงขึ้น ทำให้ดินถล่มเร็วขึ้น ซึ่งเป็นไปตามหลักวิศวกรรมปฐพีเมื่อความอิ่มตัวของดิน (saturation) เพิ่มขึ้น จะเกิดการลดแรงเฉือนของดิน (shear strength) โดยที่แรงดันน้ำในรูพรุน (pore water pressure) ในกรณีมวลดินอิ่มตัวใช้สมการที่ 6

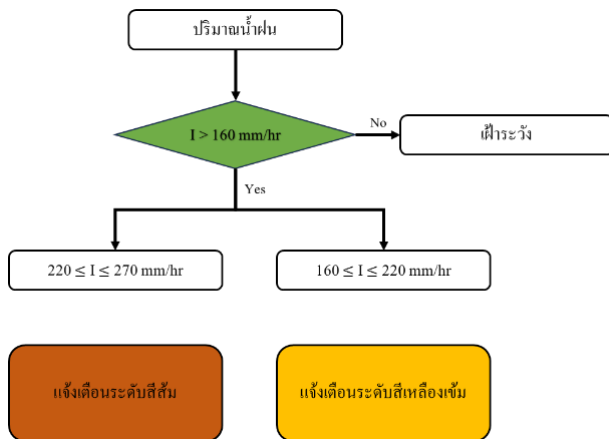
$$u = g_w \cdot h \quad (6)$$

แรงดันน้ำของแบบจำลองมีค่าเท่ากับ 5.89 kPa เท่ากันทั้ง 2 แบบจำลอง แต่ค่าแรงระหว่างการไหลของน้ำ (seepage force) ภายในมวลดินที่มีหน่วยน้ำหนักต่างกันทำให้พฤติกรรมดินถล่มของแบบจำลองแตกต่างกัน ซึ่งมีความสำคัญมากในการอธิบายกระบวนการที่ทำให้ดินสูญเสียเสถียรภาพและนำไปสู่ดินถล่ม โดยที่แรงซึมน้ำเป็นผลมาจากการไหลของน้ำผ่านรูพรุนในมวลดินมีแรงกระทำต่อเม็ดดินโดยตรงและอาจช่วยผลักดันให้หลุดจากกันได้ มวลดินที่มีหน่วยน้ำหนัก 15.30 kN/m³ จะมีอัตราส่วนแรงการไหลน้ำต่อน้ำหนักดินเท่ากับ 64% ในขณะที่แบบจำลองมวลดินที่มีหน่วยน้ำหนัก 18.34 kN/m³ จะมีค่า 53% ทำให้แบบจำลองมวลดินเดิมข้างทางเป็นดินหลวมมีแนวโน้มจะสูญเสียแรงยึดเหนี่ยวและเกิดดินถล่มสูงกว่า เนื่องจากความต้านทานมีน้อยกว่า ในขณะที่ดินบดอัดแม้มีความต้านทานสูงหากยังมีปริมาณฝนตกที่มากขึ้นจะยังคงมีแนวโน้มลดลงอย่างต่อเนื่อง เมื่อถึงจุดวิกฤตแล้วก็จะเกิดการถล่มอย่างรวดเร็ว [29] นอกจากปริมาณน้ำฝนและหน่วยน้ำหนักดินมีผลต่อระยะเวลาและความรุนแรงของดินถล่มแล้ว ลักษณะพฤติกรรมของมวลดินก่อนเกิดดินถล่มยังแตกต่างกันอีกด้วย

ในการวิเคราะห์หาระดับความรุนแรง (R) ของอัตราดินถล่มโดยใช้ปัจจัยของอัตราดินถล่ม เวลาการเกิดดินถล่มและปริมาณน้ำฝน โดยวิเคราะห์ได้ว่า ค่า (R) จะอยู่ระหว่าง 0-1 โดยสอดคล้องกับปริมาณของฝนที่ I : 100-270 มิลลิเมตรต่อชั่วโมงตามลำดับ จึงตั้งระดับการเตือนเป็นสีจากปริมาณน้ำฝนที่สอดคล้องกัน ดังรูปที่ 12 สำหรับดินเดิมข้างทางที่มีหน่วยน้ำหนัก 15.30 kN/m³ และรูปที่ 13 สำหรับดินบดอัดที่ความหนาแน่น 18.34 kN/m³ โดยระดับการเตือนล่วงหน้าและลักษณะดินถล่มที่ได้จากการศึกษาของแบบจำลอง



รูปที่ 12 : ระบบเตือนภัยล่วงหน้าดินถล่มของมวลดินเดิมข้างทางที่มีหน่วยน้ำหนัก 15.30 kN/m³



รูปที่ 13 : ระบบเตือนภัยล่วงหน้าดินถล่มของมวลดินบดอัดไหล่ทาง
ที่หน่วยน้ำหนักดิน 18.34 kN/m³

นอกจากนี้การจำแนกลักษณะพฤติกรรมการเกิดดินถล่มที่ได้จากแบบจำลอง สามารถนำข้อมูลมาจำแนกเพื่อสร้างตารางดัชนีระบบเตือนภัยดินถล่มเนื่องจากปริมาณน้ำฝนได้ ค่า R (อัตราความรุนแรงของดินถล่ม) ตั้งแต่ 0–1 โดยที่มวลดินเดิมข้างทางที่มีหน่วยน้ำหนัก 15.30 kN/m³ ควรเฝ้าระวังเมื่อมีระดับฝนตกตั้งแต่ 60 มิลลิเมตรต่อชั่วโมง และมวลดินบดอัดไหล่ทางที่หน่วยความหนาแน่น 18.34 kN/m³ ควรเฝ้าระวังเมื่อมีฝนตกตั้งแต่ 160 มิลลิเมตรต่อชั่วโมง เป็นต้นไป

การศึกษาพฤติกรรมของแบบจำลองมวลดินที่ใช้ได้คำนึงถึงปัจจัยที่จะทำให้เกิดดินถล่ม 2 บริเวณ ที่มักจะเกิดขึ้นบนเส้นทางหลวงหมายเลข 1096 อำเภอสะเมิง จังหวัดเชียงใหม่ ถนนถูกตัดผ่านภูเขาสูงชันและมีการออกแบบพิจารณาตามภูมิประเทศทำให้มีความลาดชันของดิน 70 องศา ดินเดิมบริเวณข้างทางที่ถูกตัดผ่านมีหน่วยน้ำหนัก 15.30 kN/m³ ควรเฝ้าระวัง (yellow alert) ให้เจ้าหน้าที่ตรวจสอบพื้นที่ หากฝนเกิน 100 มิลลิเมตรต่อชั่วโมง ระดับเตือนภัย (orange alert) เตรียมอพยพหรือปิดถนนเมื่อฝนเกิน 160 มิลลิเมตรต่อชั่วโมง และระดับวิกฤต (red alert) มากกว่า 220 มิลลิเมตรต่อชั่วโมง สั่งปิดเส้นทางทันทีและแจ้งเตือนประชาชนในพื้นที่เสี่ยง ในขณะที่ดินที่ถูกบดอัดเป็นไหล่ทาง มีหน่วยน้ำหนักวัดได้ 18.34 kN/m³ เมื่อปริมาณน้ำฝนมากกว่า 220 มิลลิเมตรต่อชั่วโมง จะเป็นระดับเตือนภัย เตรียมอพยพหรือปิดถนน เนื่องจากมีพฤติกรรมดินถล่มเป็นแบบดินถล่มอย่างฉับพลันบริเวณหน้าความลาดชัน ดินทั้ง 2 แบบ ได้จำแนกลักษณะรูปแบบของดินถล่มสัมพันธ์กับปริมาณน้ำฝน ดังตารางที่ 6 และ 7

ตารางที่ 6 : ดัชนีเตือนภัยดินถล่ม หน่วยน้ำหนักมวลดิน 15.30 kN/m³

ปริมาณน้ำฝน (มิลลิเมตรต่อชั่วโมง)	ระดับความ รุนแรง	รูปแบบของดินถล่ม
$I \geq 270$	1	ดินถล่มจะเกิดขึ้นถึง ส่วนกลาง มีอย่างต่อเนื่อง และเกิดขึ้นอย่างรวดเร็ว
$220 \leq I \leq 270$	0.76–1	เกิดดินถล่มหลายครั้งอย่างต่อเนื่อง
$160 \leq I \leq 220$	0.35–0.76	เกิดดินถล่มหลายครั้ง แต่ไม่ ต่อเนื่อง
$100 \leq I \leq 160$	0.25–0.35	ดินถล่มบริเวณด้านหน้าของ ความลาดชันที่ละเอียด ๆ
$60 \leq I \leq 100$	0–0.25	การสะสมตัวของน้ำในดินก่อน เกิดดินถล่ม

ตารางที่ 7 : ดัชนีเตือนภัยดินถล่ม หน่วยน้ำหนักมวลดิน 18.34 kN/m³

ปริมาณน้ำฝน (มิลลิเมตรต่อชั่วโมง)	ระดับความ รุนแรง	รูปแบบของดินถล่ม
$I \geq 270$	1	การพังทลายอย่างฉับพลันถึง ส่วนกลางของความลาดชัน
$220 \leq I \leq 270$	0.72–1	ดินถล่มอย่างฉับพลันบริเวณ หน้าความลาดชัน
$160 \leq I \leq 220$	0–0.72	เกิดการสะสมตัวของน้ำในดิน จนทำให้เกิดการพังทลาย บริเวณหน้าความลาดชัน
$100 \leq I \leq 160$	0	การสะสมตัวของน้ำในดินก่อน เกิดการพังทลาย

5) สรุปผล

จากการศึกษาแบบจำลองมวลดินในท้องปฏิบัติการ โดยใช้ดินที่มีหน่วยน้ำหนัก 15.30 kN/m³ สำหรับการจำลองดินเดิมบริเวณข้างทางที่ถูกถนนตัดผ่าน และดินบดอัดไหล่ทางที่มีหน่วยน้ำหนัก 18.34 kN/m³ พบว่ามีขีดจำกัดปริมาณน้ำฝนเพื่อเฝ้าระวังการเกิดดินถล่มที่ 60 และ 160 มิลลิเมตรต่อชั่วโมงตามลำดับการจัดทำระดับการแจ้งเตือน (warning levels) ควรอาศัยเครือข่ายสถานีวัดฝนอัตโนมัติที่ครอบคลุมพื้นที่ภูเขาอย่างเพียงพอเพื่อให้ได้ข้อมูลฝนแบบเรียลไทม์คล้ายกับกรณีของฮ่องกงที่มีสถานีวัดฝนกว่า 120 สถานีให้ข้อมูลแก่ระบบเตือนดินถล่มตลอดเวลา [17], [30]

การศึกษาแบบจำลองดินถล่มยังสามารถจำแนกพฤติกรรมลักษณะดินถล่มได้ 4 ลักษณะ ตั้งแต่การถล่มช้า ๆ เฉพาะขอบด้านบนจนถึงการถล่มอย่างรวดเร็วทั้งหมด ซึ่งช่วยพัฒนาเกณฑ์

เดือนภัยที่สัมพันธ์กับปริมาณฝน คือ ระดับสีเหลือง ที่ปริมาณน้ำฝน $>60 \leq 100$ มิลลิเมตรต่อชั่วโมง ระดับสีเหลืองเข้มที่ปริมาณน้ำฝน $>100 \leq 160$ มิลลิเมตรต่อชั่วโมง ระดับสีส้มปริมาณน้ำฝน $>160 \leq 220$ มิลลิเมตรต่อชั่วโมง และระดับสีแดงปริมาณน้ำฝน >220 มิลลิเมตรต่อชั่วโมง สำหรับมวลดินหลวมข้างทางซึ่งมักจะเกิดดินถล่มลงมาปิดทับเส้นทางในพื้นที่กรณีศึกษา

ผลการศึกษาสอดคล้องกับงานวิจัยที่พบว่าปริมาณน้ำฝนระดับ 30–60 มิลลิเมตรต่อชั่วโมงส่งผลต่อการถล่มของมวลดินได้ และสัดส่วนแรงซึมผ่านของน้ำและน้ำหนักของดินทั้ง 2 หน่วยน้ำหนักมีมากกว่าร้อยละ 50% เนื่องจากการสะสมความชื้นในดินตามระยะเวลาส่งผลต่อการถล่มมวลดิน [15], [18] นำไปสู่การวางแผนจัดระบายน้ำและกำแพงกันดินให้เหมาะสมกับปริมาณฝนที่คาดการณ์ว่าจะเกิดขึ้นได้ เพื่อป้องกันน้ำฝนสะสมจนทำให้ลาดดินอึดตัวเกินกว่าจุดวิกฤต นอกจากนี้การจัดการพื้นที่ต้นน้ำ เช่น การปลูกป่าและรักษาหน้าดิน ก็เป็นส่วนหนึ่งของแผนระยะยาวที่ช่วยลดความรุนแรงของน้ำหลากและชะลอการซึมของน้ำลงดิน ลดโอกาสเกิดดินถล่มในพื้นที่ลาดชัน

6) ข้อเสนอแนะ

หน่วยงานรัฐควรเร่งติดตั้งระบบวัดฝนและเซ็นเซอร์ตรวจวัดการเคลื่อนตัวของดิน พร้อมจัดสรรงบประมาณสำหรับงานป้องกันเชิงโครงสร้าง และการเสริมความเข้าใจแก่ประชาชนในพื้นที่ นอกจากนี้การศึกษาเพิ่มเติมเพื่อพัฒนาแบบจำลองผสมผสานทั้งแบบกายภาพและด้วยโปรแกรมคอมพิวเตอร์ โดยใช้ AI วิเคราะห์ข้อมูลฝนและใช้เซ็นเซอร์จากภาคสนาม พร้อมเชื่อมโยงกับเทคโนโลยี Remote Sensing เช่น InSAR หรือ LiDAR เพื่อเฝ้าระวังแบบพื้นที่ขนาดใหญ่ และทดลองจำลองการถล่มในขนาดที่ใหญ่ขึ้นเพื่อทดสอบสมมติฐานเชิงพลศาสตร์ของดิน [16], [31]

ทั้งนี้การศึกษาในห้องปฏิบัติการสามารถควบคุมตัวแปรได้ดี แต่ก็มีข้อจำกัดด้านการรบกวนและความซับซ้อนของภูมิประเทศจริง [32] โดยการศึกษาเน้นจำลองฝนระยะสั้นในขนาดจำกัด ซึ่งอาจไม่ครอบคลุมสภาพแวดล้อมจริงที่มีปัจจัยร่วมหลายด้าน เช่น ฝนสะสมหลายวัน รากพืช หรือแรงกระทำจากแผ่นดินไหว [13], [14] งานวิจัยนี้แสดงให้เห็นว่าแบบจำลองทางกายภาพสามารถให้ข้อมูลเชิงกลไกที่นำไปประยุกต์ใช้ได้จริงในการพัฒนาเกณฑ์เตือนภัยดินถล่ม เกณฑ์ฝนที่ได้ (60 และ 160 มิลลิเมตรต่อชั่วโมง) สามารถนำไปใช้ในการบริหารจัดการพื้นที่เสี่ยงและการ

วางแผนเชิงนโยบายได้อย่างมีประสิทธิภาพ หากมีการศึกษาและการสนับสนุนด้วยข้อมูลสนาม เทคโนโลยี และบูรณาการข้ามหน่วยงาน

REFERENCES

- [1] F. Guzzetti, A. C. Mondini, M. Cardinali, F. Fiorucci, M. Santangelo, and K. T. Chang, "Landslide inventory maps: new tools for an old problem," *Earth-Sci. Rev.*, vol. 112, no. 1–2, pp. 42–66, 2012.
- [2] R. Gracheva and A. Golyeva, "Landslides in mountain regions: hazards, resources and information," in *Geophysical Hazards*, T. Beer, Ed. Dordrecht, Netherlands: Springer, 2010, pp. 249–260.
- [3] J. Benoit *et al.*, "The Oso, Washington, USA landslide: Lessons learned through field reconnaissance and data collection," in *From Fundamentals to Applications in Geotechnics*, Amsterdam, Netherlands: IOS Press, 2016, pp. 3167–3174. doi: 10.3233/978-1-61499-603-3-3167.
- [4] K. Ho, S. Lacasse, and L. Picarelli, Eds., *Slope Safety Preparedness for Impact of Climate Change*. Boca Raton, FL, USA: CRC Press, 2017.
- [5] R. A. Yuniawan *et al.*, "Revised rainfall threshold in the Indonesian landslide early warning system," *Geosciences*, vol. 12, 2022, Art. no. 129, doi: 10.3390/geosciences12030129.
- [6] K. Intarat, P. Yoomee, A. Hussadin, and W. Lamprom, "Assessment of landslide susceptibility in the intermontane Basin Area of Northern Thailand," *Environ. Nat. Resour. J.*, vol. 22, no. 2, pp. 158–170, 2024.
- [7] C. Dechkamfoo, *et al.*, "Impact of rainfall-induced landslide susceptibility risk on mountain roadside in Northern Thailand," *Infrastructures*, vol. 7, no. 2, 2022, Art. no. 17, doi: 10.3390/infrastructures7020017.
- [8] Naewna. "Avoid the route! 'Chiang Mai' landslide blocks traffic," (in Thai), NAEWNANEWS.com. <https://www.naewna.com/local/827336> (accessed Sep. 6, 2024).
- [9] DPM Reporter, "A landslide has blocked Highway 1096," (in Thai), 2022. [Online]. Available: <https://dpmreporter.dhs.go.th/portal/disaster-news/3659>
- [10] Chiang Mai News. "Landslide blocks Mae Rim–Samoeng road near the botanical garden," (in Thai), CHIANGMAI NEWS.co.th. <https://www.chiangmainews.co.th/disaster/2731955/> (accessed Sep. 13, 2022).



- [11] P. Yongsiri, P. Soyotong, W. Laongmanee, and P. Uthaisiri, "Development of a landslide risk area display system by using the geospatial technology with daily rainfall data via the internet network in the Northern region of Thailand," *Int. J. Agricultural Technol.*, vol. 19, no. 4, pp. 1953–1968, 2023.
- [12] S. Zhang et al., "Model test study on the hydrological mechanisms and early warning thresholds for loess fill slope failure induced by rainfall," *Eng. Geol.*, vol. 258, 2019, Art. no. 105135.
- [13] S. Phumcha-em, B. Kunsuwan, and W. Mairiang, "3-D seepage modeling for analyzing unsaturated slope stability," *Naresuan Univ. Eng. J.*, vol. 12, no. 2, pp. 73–84, 2017.
- [14] N. Casagli et al., "Monitoring and early warning systems: Applications and perspectives," in *Understanding and Reducing Landslide Disaster Risk*, B. Šarić, K. Sassa, P. Canuti, and Y. Yin, Eds. Cham, Switzerland: Springer, 2021, pp. 1–21.
- [15] Q. Li, D. Huang, S. Pei, J. Qiao, and M. Wang, "Using physical model experiments for hazards assessment of rainfall-induced debris landslides," *J. Earth Sci.*, vol. 32, pp. 1113–1128, 2021.
- [16] R. I. Spiekermann, F. V. Zadelhoff, J. Schindler, H. Smith, C. Phillips, and M. Schwarz, "Comparing physical and statistical landslide susceptibility models at the scale of individual trees," *Geomorphology*, vol. 440, 2023, Art. no. 108870.
- [17] Y. Qin, G. Yang, K. Lu, Q. Sun, J. Xie, and Y. Wu "Performance evaluation of five GIS-based models for landslide susceptibility prediction and mapping: A case study of Kaiyang county, China," *Sustainability*, vol. 13, no. 11, 2021, Art. no. 6441.
- [18] P. Sun, H. Wang, G. Wang, R. Li, Z. Zhang, and X. Huo, "Field model experiments and numerical analysis of rainfall-induced shallow loess landslides," *Eng. Geo.*, vol. 295, Dec. 2021, Art. no. 106411.
- [19] S. Soralump, "Rainfall-Triggered Landslide: from research to mitigation practice in Thailand", *Geotechnical Eng.*, vol. 41, no. 1, pp. 39–44, 2024.
- [20] A. Chinkulkijniwat, R. Salee, S. Horpibulsuk, A. Arulrajah, and M. Hoy, "Landslide rainfall threshold for landslide warning in Northern Thailand," *Geomatics, Natural Hazards and Risk*, vol. 13, no. 1, pp. 2425–2441, 2022.
- [21] J. Sukpinit, P. Hemwan, and A. Charoenpanyanet, "Landslide susceptibility assessment using frequency ratio method: A case study in Sakad village, Sakad Subdistrict, Pua District, Nan Province," *Burapha Sci. J.*, vol. 27, no. 3, pp. 1832–1851, 2022.
- [22] Weather Spark. "Climate and average weather year round in Samoeng" [Online]. Available: <https://weatherspark.com/y/112815/Average-Weather-in-Samoeng-Thailand-Year-Round#Figures-Rainfall> (accessed Mar. 7, 2024).
- [24] Department of Mineral Resources, Ministry of Natural Resources and Environment, *Zoning for Geological Management and Mineral Resources of Chiang Mai Province*, 2nd ed. Bangkok, Thailand: Amarin Printing & Publishing (in Thai), 2015.
- [25] D. C. Wyllie and C. W. Mah, *Rock Slope Engineering: Civil and Mining*, 4th ed., New York, NY, USA: Spon Press, 2004.
- [26] Geospatial JS. "Principles of slope design for road and grading works (basic)," BLOGSPOT.com. <https://geomatics-tech.blogspot.com/2013/10/slope-grading.html> (accessed Jul. 1, 2024)
- [27] J.-C. Chen and W.-S. Huang, "Evaluation of rainfall-triggered debris flows under the impact of extreme events: A Chenyulan watershed case study, Taiwan," *Water*, vol. 13, no. 16, 2021, Art. no. 2201.
- [28] Y.-m. Wu, H.-x. Lan, X. Gao, L.-p. Li and Z.-h. Yang, "A simplified physically based coupled rainfall threshold model for triggering landslides," *Eng. Geo.*, vol. 195, pp. 63–69, 2015.
- [29] J. A. Hudson and J. P. Harrison, *Engineering Rock Mechanics: An Introduction to the Principles*, New York, NY, USA: Elsevier Science, 1997.
- [30] S. Segoni, L. Piciullo, and S. L. Gariano, "Preface: Landslide early warning systems: Monitoring systems, rainfall thresholds, warning models, performance evaluation and risk perception," *Nat. Hazards Earth Syst. Sci.*, vol. 18, pp. 3179–3186, 2018.
- [31] P. Khiaosalap and P. Tongdenok, "Assessment of landslide hazard applying techniques weighted factor index method and geographic information system in Huai Mae Saroi watershed, Phrae province," in *Proc. 53rd Kasetsart Univ. Annu. Conf.: Sci. Genetic Eng., Architecture and Eng., Agro-Industry, Natural Resour. and Environ.*, Feb. 2015, pp. 1264–1271.
- [32] L. Piciullo, M. Calvello and J. M. Cepeda, "Territorial early warning systems for rainfall-induced landslides," *Earth-Sci. Rev.*, vol. 179, pp. 228–247, 2018.

การแก้ปัญหาการบำรุงรักษาของโรงงาน กรณีศึกษาของบริษัทผลิต และจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์

อดุลย์ พุกอินทร์*

*หลักสูตรเทคโนโลยีอุตสาหกรรม คณะเทคโนโลยีอุตสาหกรรม มหาวิทยาลัยราชภัฏอุดรดิตถ์, อุดรดิตถ์, ประเทศไทย

*ผู้ประพันธ์บรรณกิจ อีเมล : adun999@gmail.com

รับต้นฉบับ : 16 สิงหาคม 2567; รับประทานฉบับแก้ไข : 28 ตุลาคม 2567; ตอบรับบทความ : 3 ธันวาคม 2567

เผยแพร่ออนไลน์ : 26 ธันวาคม 2568

บทคัดย่อ

การวิจัยนี้มุ่งเน้นการจัดตารางการบำรุงรักษาของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ โดยได้พัฒนาแบบจำลองกำหนดการเชิงคณิตศาสตร์ร่วมกับการใช้อัลกอริทึมฮิวริสติกส์โครงข่ายประสาทเทียม (Neural Network: NNs) และวิธีการโลคอลเสิร์ช (Local Search: LS) เพื่อแก้ปัญหาการบำรุงรักษาแบบ Preventive Maintenance (PM) ซึ่งจำเป็นสำหรับการผลิตที่ต้องดำเนินการอย่างต่อเนื่อง การวิจัยได้ทดสอบโปรแกรมที่พัฒนาขึ้นกับการทดลองเชิงแฟกทอเรียล 2^3 เพื่อหาค่าพารามิเตอร์ที่เหมาะสมของการค้นหาคำตอบ การวิจัยได้เก็บข้อมูลปัญหาขนาดเล็กและขนาดใหญ่ ซึ่งโปรแกรมมีความสามารถหาค่าคำตอบในการจัดตารางการบำรุงรักษาได้อย่างมีประสิทธิภาพ โดยได้ค่าเมคสเปน (Makespan) ที่ใกล้เคียงและตรงกับค่าขอบเขตต่ำสุด (Lower Bound) ซึ่งสะท้อนถึงประสิทธิภาพของโครงข่ายประสาทเทียม และวิธีการโลคอลเสิร์ชในการแก้ปัญหา นอกจากนี้ การเก็บข้อมูลก่อนและหลังการวิจัยในช่วง 6 เดือน พบว่า ต้นทุนรวมลดลงเหลือ 1,915,062 บาท ลดลงจากเดิมเท่ากับ 422,396 บาท คิดเป็นร้อยละ 9.93 ส่วนค่าระยะเวลาระหว่างการเกิดเหตุขัดข้องของเครื่องจักร (MTBF) มีค่าเพิ่มขึ้นเป็น 83.84 ชั่วโมง คิดเป็นร้อยละ 35.27 แสดงให้เห็นถึงการลดลงของต้นทุนทั้งในด้านเวลาและการบำรุงรักษา

คำสำคัญ : ค่าขอบเขตต่ำสุด แบบจำลองกำหนดการเชิงคณิตศาสตร์ การเกิดเหตุขัดข้องของเครื่องจักร ผลิตภัณฑ์อาหารกึ่งสำเร็จรูป



Solving Factory Maintenance Problem: A Case Study of a Semi-Finished Food Product Manufacturing and Distribution Company in Uttaradit Province

Adun Phuk-in*

**Industrial Technology Program, Faculty of Industrial Technology, Uttaradit Rajabhat University, Uttaradit, Thailand*

*Corresponding Author. E-mail address: adun999@gmail.com

Received: 16 August 2024; Revised: 28 October 2024; Accepted: 3 December 2024

Published online: 26 December 2025

Abstract

This research focuses on scheduling maintenance for a semi-finished food product manufacturing and distribution company in Uttaradit Province. Heuristic algorithms, neural networks (NNs), and local search (LS) were used to create a mathematical scheduling model that solves the preventive maintenance (PM) problem. This is needed for production to keep going. The researchers tested the developed program with 2^3 factorial experiments to find the appropriate parameter values for the answer. The research collected data on both small and large problems, and the program was able to find the answer value for scheduling maintenance efficiently. It obtained a makespan value, which was close to and matched the lower bound, reflecting the efficiency of the neural network. The local search method was employed to solve the problem. In addition, the data collected before and after the research for 6 months found that the total cost decreased to 1,915,062 baht, down from the original 422,396 baht, or 9.93 percent. The mean time between machine failures (MTBF) increased to 83.84 hours or 35.27 percent, showing a decrease in costs in terms of time and maintenance.

Keywords: Lower bound, Mathematical scheduling model, Mean Time Between Failures (MTBF), Semi-finished food products

1) บทนำ

การแปรสภาพวัตถุดิบให้เป็นผลิตภัณฑ์สำเร็จรูปในกระบวนการผลิต [1] จำเป็นต้องใช้เครื่องจักรให้เกิดประสิทธิภาพสูงสุด โดยสามารถทำได้ด้วยการลดระยะเวลาการทำงาน การลดขั้นตอนหรือวิธีการ การลดต้นทุนในการใช้พลังงาน การลดต้นทุนการขนส่ง และการเพิ่มประสิทธิภาพของพนักงานในการใช้เครื่องจักร เป็นต้น ประสิทธิภาพของเครื่องจักรส่งผลต่อต้นทุนการผลิต ซึ่งขึ้นอยู่กับการจัดการขั้นตอนหรือวิธีการที่เหมาะสมในการใช้เครื่องมือ หรือการนำวิธีการมาใช้ในการแก้ไขปัญหาในรูปแบบต่าง ๆ ตามกระบวนการออกแบบเชิงวิศวกรรม [2], [3] ซึ่งมีการดำเนินงาน 6 ขั้นตอน [4]–[6] คือ 1) ระบุปัญหา (problem identification): การระบุและกำหนดปัญหาที่ต้องการแก้ไขในกระบวนการผลิต 2) รวบรวมข้อมูลและแนวคิดที่เกี่ยวข้องกับปัญหา (related information search): การเก็บรวบรวมข้อมูลและแนวคิดที่เกี่ยวข้องกับปัญหาเพื่อใช้ในการวิเคราะห์ 3) ออกแบบวิธีการแก้ปัญหา (solution design): การออกแบบและพัฒนารูปแบบการแก้ปัญหาที่เหมาะสม 4) วางแผนและดำเนินการแก้ปัญหา (planning and development): การวางแผนและดำเนินการแก้ปัญหาตามวิธีการที่ออกแบบ 5) ทดสอบ ประเมินผลและปรับปรุงวิธีการแก้ปัญหา (testing, evaluation, and design improvement): การทดสอบและประเมินผลวิธีการแก้ปัญหา และการปรับปรุงตามผลที่ได้ 6) นำเสนอวิธีการแก้ปัญหา ผลการแก้ปัญหา (presentation): การนำเสนอวิธีการแก้ปัญหา

ผลการออกแบบเชิงวิศวกรรมจึงมี ขั้นตอนประเมินและการปรับปรุงกระบวนการ (process evaluation and improvement) การประเมินกระบวนการทั้งหมดเพื่อหาจุดที่สามารถปรับปรุงได้เพิ่มเติมเพื่อเพิ่มประสิทธิภาพและลดต้นทุน และขั้นตอนตรวจสอบและควบคุมคุณภาพ (quality control and assurance) การตรวจสอบและควบคุมคุณภาพของผลิตภัณฑ์สำเร็จรูปเพื่อให้แน่ใจว่าผลิตภัณฑ์มีคุณภาพตามมาตรฐานที่กำหนด เพิ่มเข้ามาทำให้เกิดความชัดเจนได้มากยิ่งขึ้น [1], [6]

การรักษาประสิทธิภาพสูงสุดในกระบวนการผลิตเป็นสิ่งสำคัญ เพื่อเพิ่มความสามารถในการแข่งขันในตลาด ปัญหาหลักที่ส่งผลต่อประสิทธิภาพการผลิต [7] ได้แก่ การใช้งานเครื่องจักรไม่เต็มประสิทธิภาพ ซึ่งอาจเกิดจากการขาดการบำรุงรักษาที่เหมาะสม การตั้งค่าเครื่องจักรไม่ถูกต้อง หรือการใช้งานไม่ต่อเนื่อง นอกจากนี้ยังมีปัญหาเรื่องการขาดการวางแผนการผลิตที่มีประสิทธิภาพ

การหยุดชะงักที่ไม่จำเป็น การจัดการคุณภาพไม่ดี และการส่งวัตถุดิบล่าช้า ปัญหาเหล่านี้ลดประสิทธิภาพการผลิตและเพิ่มต้นทุน [8], [9] การแก้ไขต้องอาศัยการวางแผนและการจัดการการผลิตที่ดี การใช้เทคโนโลยีทันสมัย และการพัฒนาบุคลากรอย่างต่อเนื่อง เพื่อให้ทรัพยากรให้เต็มประสิทธิภาพและลดต้นทุนการผลิตที่เป็นสิ่งจำเป็น [10]–[12]

จากการศึกษาข้อมูลการผลิตของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรธานี ซึ่งมีการผลิตอาหารจากข้าวและแป้งมันสำปะหลัง มาเป็นผลิตภัณฑ์เส้นก๋วยเตี๋ยว เส้นก๋วยจั๊บ เส้นใหญ่ เส้นบะหมี่ เบเกอรี่ ขนมปัง และอื่น ๆ ที่เป็นผลิตภัณฑ์จากแป้ง เป็นโรงงานที่มีกำลังการผลิตสูง และมียอดขายมากที่สุดในภาคเหนือ มีการส่งจำหน่ายไปยังจังหวัดต่าง ๆ เช่น เชียงใหม่ เชียงราย ลำปาง น่านแพร่ พิชญ์โลก และในประเทศไทย เป็นต้น กระบวนการผลิตใช้วัตถุดิบประเภทแป้งข้าวเจ้า แป้งมันสำปะหลัง และส่วนผสมตามสูตรของโรงงาน มีกำลังการผลิตมากกว่า 1.5 ตันต่อวัน มียอดขายมากกว่า 90 ล้านบาทต่อปี มีเครื่องจักรที่ใช้ในการผลิตแบบต่อเนื่องทั้งแบบอัตโนมัติ และการทำงานแบบกึ่งอัตโนมัติ (semi-automatic work) เช่น หม้อต้ม (boiler) หม้อนึ่งไอน้ำ เครื่องอบความร้อนระบบอัตโนมัติ และเครื่องตัดเส้นระบบอัตโนมัติ เป็นต้น เครื่องจักรโรงงานนี้มีจำนวนมาก การวิจัยจึงนำปัญหาของการบำรุงรักษามาใช้ในการจัดการการบำรุงรักษา ซึ่งในปัจจุบันเป็นแบบซ่อมแซมหลังเกิดเหตุขัดข้อง (breakdown maintenance) หรือการใช้การพยากรณ์ในบางเครื่องจักร ซึ่งพบปัญหาการรอคอยอะไหล่ในการซ่อม หรือไม่ได้มีการวางแผนการจัดเก็บอะไหล่ในการซ่อมที่เกิดความผิดพลาด ทำให้เกิดต้นทุนการผลิตทางตรงและทางอ้อม

การวิจัยนี้จึงมีแนวทางในการพัฒนาแบบจำลองกำหนดการเชิงคณิตศาสตร์ เพื่อใช้ในการจัดการการบำรุงรักษา และการพัฒนาอัลกอริทึมโครงข่ายประสาทเทียม (Neural Network: NNs) ร่วมกับวิธีการโลคอลเสิร์ช (Local Search: LS) ที่จะนำมาใช้แก้ปัญหาการวางแผนการจัดการการบำรุงรักษา เพื่อหาค่าเมคสแปน (Makespan) ที่มีค่าน้อยที่สุด และเพื่อเปรียบเทียบค่าระยะเวลาเฉลี่ยระหว่างการเกิดเหตุขัดข้องของเครื่องจักร (Mean Time Between Failure: MTBF) ค่าระยะเวลาเฉลี่ยระหว่างการเกิดเหตุขัดข้องของเครื่องจักร และการวัดผลประสิทธิภาพจากร้อยละของเวลาที่เครื่องจักรเกิดเหตุขัดข้อง (Average Machine Downtime: AMD) กับการผลิตในปี พ.ศ. 2566 และปี พ.ศ. 2567 จำนวน 12 เดือน

2) ทบทวนวรรณกรรม

2.1) ปัญหาการจัดตารางงาน (Scheduling Problem) [11]

เป็นปัญหาสำคัญในอุตสาหกรรมการผลิต ซึ่งต้องมีการวางแผนการผลิตที่มีประสิทธิภาพในการจัดตารางงานของเครื่องจักรภายใต้ข้อบังคับและข้อจำกัดต่าง ๆ เช่น เวลา เครื่องจักร และทรัพยากรบุคคล [12] การจัดตารางงานที่เหมาะสมจะช่วยให้การผลิตดำเนินไปอย่างราบรื่น และตรงตามกำหนดการส่งมอบให้กับลูกค้า [13], [14] การวิจัยนี้จะนำวิธีการจัดตารางงานมาใช้ในการวางแผนการจัดตารางการบำรุงรักษาเครื่องจักรในการผลิตเพื่อลดต้นทุนการบำรุงรักษา ซึ่งเป็นประเด็นสำคัญในการวิจัยนี้

2.2) การบำรุงรักษาเชิงป้องกัน (Preventive Maintenance: PM) [15]

หนึ่งในรูปแบบการบำรุงรักษาที่ได้รับความนิยม เช่นเดียวกับการบำรุงรักษาที่ทุกคนมีส่วนร่วม (TPM) [16] โรงงานของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรธานี เลือกใช้การบำรุงรักษานี้ เนื่องจากโรงงานเป็นขนาดกลาง ซึ่งรูปแบบการบริหารจัดการเน้นการดูแลสภาพเครื่องจักรและอุปกรณ์ภายในโรงงาน ที่ใช้การตรวจสอบ ซ่อมแซม หรือเปลี่ยนแปลงอุปกรณ์ต่าง ๆ ตามเวลาที่มีการกำหนดไว้ตามแผนการบำรุงรักษา [8] โดยเป็นการดำเนินการที่มุ่งเน้นการตรวจสอบและบำรุงรักษาเครื่องจักร [17], [18] หรืออุปกรณ์อย่างสม่ำเสมอ ก่อนที่ปัญหาจะเกิดขึ้น โดยมีวัตถุประสงค์หลัก เพื่อป้องกันการเสียหายหรือความล้มเหลวที่อาจเกิดขึ้น จากการหยุดชะงักในการผลิต และยืดอายุการใช้งานของเครื่องจักร

การบำรุงรักษาเชิงป้องกัน มีหลักการบำรุงรักษาตามเวลา (time-based maintenance) การดำเนินการบำรุงรักษาตามตารางเวลาที่กำหนดไว้ เช่น การตรวจสอบและเปลี่ยนชิ้นส่วนตามระยะเวลาที่กำหนด และการบำรุงรักษาตามสภาพ (condition-based maintenance) การดำเนินการบำรุงรักษาตามสภาพจริงของเครื่องจักร โดยการตรวจสอบและวิเคราะห์สภาพการทำงานของเครื่องจักรเพื่อหาสัญญาณการเสื่อมสภาพหรือปัญหาที่อาจเกิดขึ้นในการผลิต ลดการเกิดความล้มเหลวของเครื่องจักร ยืดอายุการใช้งานของเครื่องจักร ลดค่าใช้จ่ายในการซ่อมแซมแบบฉุกเฉิน เพิ่มความปลอดภัยในการทำงาน และเพิ่มประสิทธิภาพการผลิต

2.3) วิธีการทางฮิวริสติกส์ (Heuristic)

เทคนิคการค้นหาที่ใช้กฎเกณฑ์ เพื่อหาโซลูชันที่ดีที่สุดสำหรับปัญหาซับซ้อน [19] โดยไม่จำเป็นต้องหาคำตอบที่ดีที่สุด เหมาะสำหรับการแก้ปัญหาที่มีข้อจำกัดด้านเวลาและทรัพยากร เช่น การหาคำตอบในปัญหาการจัดตารางงานหรือการวางแผนที่ซับซ้อน ซึ่งไม่สามารถหาคำตอบที่สมบูรณ์แบบได้ในเวลาที่เหมาะสม วิธีการทางฮิวริสติกส์ที่ใช้อยู่มีหลายวิธี [18] เช่น Greedy Algorithm, Local Search, Simulated Annealing, Genetic Algorithm, Bat Algorithm (BA) และ Tabu Search เป็นต้น และมีงานวิจัยที่นำวิธีฮิวริสติกส์มาใช้ [20], [21] เช่น การวิจัยของ A. Phuk-in [21] ได้นำปัญหาของบริษัทผลิตแป้งมันสำปะหลังของบริษัทผลิตแป้งมันสำปะหลัง ซึ่งพบปัญหาด้านการวางแผนการผลิตรวมและการจัดตารางการบำรุงรักษา การวิจัยจึงพัฒนาแบบจำลองเชิงคณิตศาสตร์ในการแก้ปัญหาพร้อมกับการพัฒนาโปรแกรมวิธีฮิวริสติกส์เจเนติกส์อัลกอริทึม (GA) การวิจัย พบว่า ในการจัดตารางงานทำให้ได้ค่าเมคสแปนต่ำสุดโดยวัดจากค่าช่องว่างการเปรียบเทียบ (Gap) ที่ดี ร่วมกับการวัดค่าการเปรียบเทียบค่าร้อยละของเวลาที่เครื่องจักรเกิดเหตุขัดข้องก่อนและหลังการวิจัย

การวิจัยของ C. Muangdit and K. Lurang [22] ได้นำเสนอการออกแบบตัวควบคุม PIDA ที่เหมาะสมที่สุดสำหรับระบบควบคุมอุณหภูมิเตาไฟฟ้า โดยใช้ Bat Algorithm (BA) ซึ่งเป็นเทคนิคการเพิ่มประสิทธิภาพแบบฮิวริสติกส์ ผลการวิจัยพบว่า PIDA ที่ออกแบบด้วย BA มีประสิทธิภาพในการควบคุมและติดตามอุณหภูมิของเตาไฟฟ้าได้ดีกว่า PIDA แบบทั่วไปได้ดี

3) วิธีดำเนินการวิจัย

3.1) การสำรวจข้อมูลของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรธานี

กระบวนการผลิตของโรงงานจากการสำรวจการผลิต และการบำรุงรักษาเครื่องจักรในปี พ.ศ. 2566 เป็นการผลิตแบบต่อเนื่อง (continuous process) ใช้วัตถุดิบในการผลิต คือ ข้าวเจ้า แป้งมันสำปะหลัง ที่มีอยู่ในพื้นที่เป็นจำนวนมาก กระบวนการผลิตจะมีการผลิตทุกวันในช่วงเวลา 08:00–18:00 น. และเป็นการผลิตทั้ง 12 เดือน ก่อนการวิจัยสำรวจข้อมูลดังนี้

3.1.1) การผลิตใช้แป้งข้าวเจ้า เป็นวัตถุดิบหลัก มีปริมาณการใช้มากกว่า 8–20 ตันต่อเดือน มีคลังสินค้าในการจัดเก็บวัตถุดิบและของคงคลังไม่มาก เนื่องจากอยู่ใกล้โรงงานและแหล่งวัตถุดิบ

ของผู้ร่วมค้า (supplier) ที่มีอยู่ในพื้นที่จังหวัดอุดรดิตถ์จำนวน 2-5 บริษัท ปริมาณการใช้มากกว่า 170 ตันต่อปี

3.1.2) การผลิตใช้ส่วนผสมแป้งมันสำปะหลังเป็นส่วนผสมตามสูตร ในการผลิตผลิตภัณฑ์อาหารของบริษัท ซึ่งแป้งมันเป็นวัตถุดิบรองที่ใช้ในการผสมกับแป้งข้าวเจ้า คือ จะมีส่วนผสมตามสูตรในแต่ละผลิตภัณฑ์ของบริษัท มีปริมาณการใช้มากกว่า 5 ตันต่อปี

3.1.3) การวางแผนการผลิตผลิตภัณฑ์ ใช้การพยากรณ์ร่วมกับการใช้ความชำนาญของผู้บริหารฝ่ายจำหน่าย และฝ่ายผลิต ใช้ประสบการณ์ตรงในการวางแผนการผลิต ซึ่งในปี พ.ศ. 2566 การเก็บข้อมูลเดือนมกราคม-ธันวาคม พ.ศ. 2566 มีการผลิตผลิตภัณฑ์สุตงตลาดมากกว่า 508 ตัน

3.1.4) การใช้เครื่องจักรในการผลิต ใช้เครื่องจักรที่เป็นเทคโนโลยีการผลิตแบบอัตโนมัติ ได้แก่ เครื่องนึ่งความร้อนขนาด 3-4 บาร์ เครื่องอบความร้อนขนาด 20 ตัน และเครื่องจักรที่ควบคุมด้วยแรงงานคนงาน เช่น เครื่องม้วนแผ่นแป้ง เครื่องหั่นเส้นเล็ก เครื่องหั่นเส้นขนาดกลาง เครื่องผลิตแผ่นก๋วยจั๊บ เครื่องบรรจุ และเครื่องซีล เป็นต้น จากการศึกษาก่อนการวิจัยไม่ได้มีการวางแผนการบำรุงรักษาที่ได้มาตรฐานมากนัก จะใช้วิธีการปรับปรุงหลังจากเกิดเหตุขัดข้อง (corrective or breakdown) ในขณะผลิต และมีต้นทุนที่เกิดขึ้น 5 ลักษณะ คือ ต้นทุนจากการแก้ไขเครื่องจักรขัดข้องในระหว่างผลิต ต้นทุนการใช้แรงงานในการซ่อมเครื่องจักร ต้นทุนในการตั้งค่าเครื่องจักร (set up time) (เนื่องจากต้องใช้คนหมู่มากในการอบ ซึ่งรอความร้อนจากหม้อต้ม) ต้นทุนการรอคอยของแรงงานในโรงงานผลิต และต้นทุนพลังงานเชื้อเพลิง การซ่อมบำรุงในปี พ.ศ. 2566 มีค่ามากกว่า 2,337,458 บาทต่อปี

3.1.5) เงินทุนหมุนเวียนที่ใช้ในการบริหารต้นทุน มีต้นทุนรวมเท่ากับ 169 ล้านบาท แบ่งเป็นต้นทุนคงที่ (fixed costs) มีเท่ากับ 32 ล้านบาท ต้นทุนผันแปร (variable costs) มีเท่ากับ 15 ล้านบาท ต้นทุนกึ่งผันแปร (semi-variable costs) มีเท่ากับ 2 ล้านบาท ต้นทุนตรง (direct costs) 102 ล้านบาท ต้นทุนทางอ้อม (indirect costs) มีเท่ากับ 5.8 ล้านบาท ต้นทุนแปรรูป (conversion costs) 3.2 ล้านบาท ต้นทุนที่เปลี่ยนแปลงตามการผลิต (operating costs) คือ ค่าซ่อมบำรุงเครื่องจักรค่าใช้จ่ายในการบำรุงรักษาอุปกรณ์ มีเท่ากับ 9 ล้านบาท

3.2) เครื่องมือที่ใช้ในการวิจัย/รวบรวมข้อมูล

การวิจัยนี้ใช้การเก็บข้อมูลภาคสนามจากโรงงาน โดยสำรวจและรวบรวมข้อมูลทั้งก่อนการวิจัยในปี พ.ศ. 2566 และหลังการวิจัยในปี พ.ศ. 2567 โดยมุ่งเน้นการวัดด้านเวลา ต้นทุน การใช้ อุปกรณ์ อะไหล่คงคลัง และต้นทุนการบำรุงรักษา กรณีศึกษาเป็นบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ งานวิจัยนำเทคโนโลยีคอมพิวเตอร์มาช่วยแก้ปัญหาที่ซับซ้อนและมีขนาดใหญ่ โดยพัฒนาโปรแกรม VBA ควบคู่กับการใช้แบบจำลองกำหนดการเชิงคณิตศาสตร์ เพื่อเพิ่มประสิทธิภาพในการแก้ปัญหา

3.3) การออกแบบการพัฒนาการวางแผนการบำรุงรักษาของบริษัท

จากการสืบค้นข้อมูลก่อนการวิจัยพบข้อมูลการบำรุงรักษามีต้นทุนที่สูง การวิจัยได้เก็บรวบรวมข้อมูล และร่วมกับฝ่ายบริหารฝ่ายผลิต และฝ่ายบำรุงรักษา ในการเลือกใช้การบำรุงรักษาเชิงป้องกัน คือการดำเนินการซ่อมบำรุงเครื่องจักรหรืออุปกรณ์อย่างเป็นระยะ โดยไม่รอให้เกิดการเสียหายก่อน โดยมีเป้าหมายเพื่อลดความเสี่ยงในการเกิดปัญหาหรือหยุดชะงักของการผลิต รวมถึงยืดอายุการใช้งานของอุปกรณ์ โดยการบำรุงรักษานี้มักจะกำหนดตามระยะเวลา หรือตามปริมาณการใช้งานที่คาดการณ์ไว้ล่วงหน้า โดยการวิจัยได้แบ่งหน่วยการบำรุงรักษาตามความเชี่ยวชาญของพนักงานเป็นสถานงาน [20], [21] การวัดผลโดยใช้ค่าระยะเวลาเฉลี่ยระหว่างการเกิดเหตุขัดข้องของเครื่องจักร [23], [24] แสดงดังสมการที่ 1, 2 และ 3

ค่าเฉลี่ยจำนวนความเสียหายของเครื่องจักร

$$\eta_{avg} = \frac{\sum_{i=1}^n n_i}{N} \quad (1)$$

ค่าระยะเวลาเฉลี่ยระหว่างการเกิดเหตุขัดข้องของเครื่องจักร

$$MTBF = \frac{T}{\eta_{avg}} \quad (2)$$

ค่าระยะเวลาเฉลี่ยระหว่างการเกิดเหตุขัดข้องของเครื่องจักร (MTBF) [20]–[24] ของบริษัท การคำนวณใช้จำนวนครั้งที่เกิดความเสียหายที่พบในช่วงเวลา T /ค่าเวลาเฉลี่ยทั้งหมดของความเสียหายของเครื่องจักรในบริษัทในแต่ละเดือน และความสัมพันธ์ระหว่าง MTTF และ MTTR โดยที่ MTTF คือ ค่าเวลาเฉลี่ยก่อนการเสียหาย (Mean Time to Failure) และ MTTR คือ ค่าเวลาเฉลี่ยในการซ่อมแซม (Mean Time to Repair)

การวัดผลประสิทธิภาพจากค่าร้อยละของเวลาที่เครื่องจักรเกิดเหตุขัดข้อง (AMD) แสดงดังสมการที่ 3

$$AMD = \frac{\%Machine\ Downtime}{Operation\ Time} \times 100 \quad (3)$$

โดยที่ เวลาของเครื่องจักรเกิดเหตุขัดข้อง (machine downtime) ในขณะที่ผลิตเป็นจำนวนครั้งที่เกิดความเสียหาย นำมาคิดคำนวณกับค่าเวลาเป็นค่าร้อยละที่ผลิตทั้งหมดของการผลิตผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในแต่ละเดือน

และการวิจัยได้พัฒนาแบบจำลองกำหนดการเชิงคณิตศาสตร์เพื่อใช้ในการจัดการการบำรุงรักษา โดยแสดงดัชนี พารามิเตอร์ ตัวแปรการตัดสินใจ ในการจัดตารางการบำรุงรักษาของบริษัท โดยกำหนดให้มีเครื่องจักรที่จะใช้ในการบำรุงรักษา n เครื่อง เครื่องจักรเหล่านี้จะต้องได้รับการบำรุงรักษาตามระยะ และเปลี่ยนถ่ายอะไหล่ตามระยะเวลาตามแผนที่ช่วงเวลา t กำหนดตัวแปรดังนี้

X_{it} = ตัวแปรตัดสินใจหากเครื่องจักร i ได้รับการบำรุงรักษาภายในระยะเวลา t จะมีค่าเป็น 1, ถ้าไม่ได้รับการบำรุงรักษาในวันที่กำหนดจะมีค่าเป็น 0

C_i = ต้นทุนการบำรุงรักษาเครื่องจักรที่ i

d_i = จำนวนวันที่ต้องบำรุงรักษาเครื่องจักรที่ i

m_t = จำนวนเครื่องจักรที่สามารถบำรุงรักษาได้ในวันที่ t

P = จำนวนวันที่เครื่องจักรแต่ละเครื่องต้องได้รับการบำรุงรักษาภายในระยะเวลา T

ฟังก์ชันวัตถุประสงค์ (objective function)

$$\text{Minimize} \quad \sum_{i=1}^n \sum_{t=1}^T C_i \cdot X_{it} \quad (4)$$

Subject to

$$\sum_{t=1}^T X_{it} = d_i \quad \forall i=1,2,3,\dots,n \quad (5)$$

$$\sum_{i=1}^n X_{it} \leq m_t \quad \forall t=1,2,3,\dots,T \quad (6)$$

$$\sum_{t=1}^T X_{it} \geq 1 \quad \forall i=1,2,3,\dots,n \quad (7)$$

$$X_{it} \in \{0,1\} \quad \forall i=1,2,3,\dots,n, \forall t=1,2,3,\dots,T \quad (8)$$

สมการที่ 4 ฟังก์ชันวัตถุประสงค์ ของการบำรุงรักษาเครื่องจักรที่ใช้ต้นทุน และการบำรุงรักษาตามแผนที่ได้ตั้งไว้มีค่าเป็นค่าต้นทุนต่ำสุด (minimize)

สมการที่ 5 แต่ละเครื่องจักรที่ใช้ในการผลิตผลิตภัณฑ์ต้องได้รับการบำรุงรักษาครบจำนวนวันที่กำหนดตามแผนการบำรุงรักษา PM

สมการที่ 6 จำนวนเครื่องจักรที่จะต้องบำรุงรักษาในแต่ละสัปดาห์ จะต้องไม่เกินความสามารถของหน่วยงานบำรุงรักษาที่กำหนดไว้

สมการที่ 7 เครื่องจักรที่ใช้ในการผลิตผลิตภัณฑ์อาหารของโรงงานต้องได้รับการบำรุงรักษาภายในระยะเวลาที่กำหนด

สมการที่ 8 ค่าของตัวแปร X_{it} จะต้องเป็นค่าไบนารี (0, 1)

3.4) การใช้วิธีการฮิวริสติกส์โครงข่ายประสาทเทียม (Neural Networks: NNs) ในการแก้ปัญหาการจัดตารางการบำรุงรักษา

เป็นวิธีที่ทันสมัยและมีประสิทธิภาพสูง เนื่องจากโครงข่ายประสาทเทียม สามารถจัดการกับข้อมูลที่มีความซับซ้อนและความไม่แน่นอนได้ดี ในกรณีของการจัดตารางการบำรุงรักษา [25] โครงข่ายประสาทเทียม จะช่วยในการตัดสินใจที่ซับซ้อนเกี่ยวกับเวลาและลำดับของงานบำรุงรักษา เพื่อให้เกิดประสิทธิภาพสูงสุดในกระบวนการบำรุงรักษา หรือการดำเนินงานของหน่วยงาน [26]–[28]

ขั้นตอนในการใช้ โครงข่ายประสาทเทียม เพื่อการจัดตารางการบำรุงรักษาของบริษัทจำหน่ายและผลิตผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์

3.4.1) การเก็บรวบรวมข้อมูล ข้อมูลที่จำเป็นต้องใช้ในการฝึก โครงข่ายประสาทเทียม [28] ได้แก่ ข้อมูลเกี่ยวกับการเสื่อมสภาพของเครื่องจักร, ระยะเวลาที่ใช้ในการบำรุงรักษา, ค่าของเวลาในการบำรุงรักษาแต่ละเครื่องจักรที่เก็บเป็นจำนวนวันที่ และข้อมูลอื่น ๆ ที่เกี่ยวข้องกับกระบวนการ การนำข้อมูลเข้า

(input) กำหนดให้ $X = [X_1, X_2, X_3, \dots, X_n]$ จะถูกป้อนเข้าโครงข่ายผ่านโปรแกรม

3.4.2) การเตรียมข้อมูล ข้อมูลที่เก็บรวบรวมต้องถูกแปลงให้อยู่ในรูปแบบที่ โครงข่ายประสาทเทียม [26], [28] สามารถเข้าใจได้ ซึ่งอาจต้องผ่านการปรับขนาด, การทำ Normalization หรือการทำ one-hot encoding ขึ้นอยู่กับประเภทของข้อมูล การจัดตารางการบำรุงรักษากำหนดให้แต่ละโหนดใช้ชั้นที่ซ่อน (hidden layers) จะรับคำสั่งอินพุตจากโหนดก่อนหน้า และคำนวณตามสมการที่ 9

$$Z_j^{(l)} = \sum_{i=1}^{n^{(l-1)}} w_{ji}^{(l)} a_i^{(l-1)} + b_j^{(l)} \quad (9)$$

โดยที่ $z_j^{(l)}$ = ค่าผลรวมถ่วงน้ำหนักของโหนด j ในชั้นที่ l
 $w_{ji}^{(l)}$ = ค่าน้ำหนักของการเชื่อมต่อระหว่างโหนด i ชั้นที่ $l-1$ และโหนด j ชั้นที่ l
 $a_i^{(l-1)}$ = ค่าที่ได้จากโหนด i ในชั้นที่ $l-1$
 $b_j^{(l)}$ = ค่าไบแอสของโหนด j ในชั้นที่ l
 $n^{(l-1)}$ = จำนวนโหนดในชั้น $l-1$

3.4.3) การออกแบบโครงสร้างของโครงข่ายประสาทเทียม การเลือกจำนวนชั้นและจำนวนโหนดในแต่ละชั้นขึ้นอยู่กับความซับซ้อนของปัญหา การใช้ Multi-Layer Perceptron (MLP) [25], [28] หรือ โครงข่ายประสาทเทียมที่ใช้สำหรับการจัดการการบำรุงรักษา ซึ่งค่าที่ได้จากสมการ $z_j^{(l)}$ จะได้นำมาผ่านฟังก์ชันกระตุ้น (activation function) เพื่อสร้างค่าที่จะส่งไปยังชั้นถัดไปของ โครงข่ายประสาทเทียม ดังสมการที่ 10

$$a_j^{(l)} = \mathfrak{S}(z_j^{(l)}) \quad (10)$$

โดยที่ $\mathfrak{S}(z_j^{(l)})$ คือ ฟังก์ชันกระตุ้น เช่น Sigmoid, ReLU, หรือ Tanh ที่ใช้กระตุ้นในการทำงานของ Network

3.4.4) การฝึกโมเดล โมเดลของโครงข่ายประสาทเทียม จะถูกฝึกด้วยข้อมูลที่เตรียมไว้ โดยใช้เทคนิคการฝึกแบบ Backpropagation และ Gradient Descent [28] เพื่อปรับน้ำหนักของโหนดในโครงข่ายของการจัดการการบำรุงรักษา ดังสมการที่ 11

$$Loss = \frac{1}{m} \sum_{i=1}^m (y^{(i)} - \hat{y}^{(i)})^2 \quad (11)$$

โดยที่ $y^{(i)}$ = ค่าผลลัพธ์จริงของข้อมูลชุดที่ i
 $\hat{y}^{(i)}$ = ค่าผลลัพธ์โมเดลคาดการณ์ในข้อมูลชุดที่ i
 m = จำนวนของข้อมูลทั้งหมดในชุดฝึก

3.4.5) การประเมินผลและปรับปรุงโมเดล หลังจากฝึกโมเดลแล้ว จะต้องทำการประเมินผลลัพธ์และปรับปรุงโมเดลให้มีความแม่นยำและประสิทธิภาพสูงขึ้น โดยใช้เทคนิคการ Cross-validation และการปรับแต่งพารามิเตอร์ (hyperparameter tuning) ซึ่งการจัดการการบำรุงรักษาออกแบบการปรับปรุง ดังสมการที่ 12

$$w_{ji}^{(l)} = w_{ji}^{(l)} - n \frac{\partial Loss}{\partial w_{ji}^{(l)}} \quad (12)$$

โดยที่ $w_{ji}^{(l)}$ คือ ค่าน้ำหนักเดิม (old weight)

N คือ ค่าอัตราการเรียนรู้ (learning rate) ซึ่งควบคุมความเร็วในการอัปเดตน้ำหนัก

$\frac{\partial Loss}{\partial w_{ji}^{(l)}}$ คือ อนุพันธ์ของค่าความสูญเสีย (loss function) เทียบกับน้ำหนัก หรือที่เรียกว่า Gradient ของน้ำหนัก

Loss คือ ฟังก์ชันการสูญเสียซึ่งวัดความแตกต่างระหว่างผลลัพธ์ที่คาดการณ์กับค่าจริง

3.4.6) การใช้งานโมเดลในสภาพแวดล้อมจริง เมื่อโมเดลถูกฝึกและปรับปรุงจนได้ผลลัพธ์ที่น่าพอใจแล้ว โมเดลสามารถนำไปใช้งานในการจัดการการบำรุงรักษาในสภาพแวดล้อมจริง โดยจะช่วยให้การตัดสินใจเป็นไปอย่างรวดเร็วและแม่นยำในชั้นเอาต์พุต $a_j^{(L)}$ จะถูกคำนวณค่าคล้ายกับชั้นที่ซ่อน และผลลัพธ์ที่ได้จะเป็นผลลัพธ์ของโครงข่ายค่าที่ดีของการจัดการการบำรุงรักษาของบริษัทรวมกับการได้ค่าเมตริกที่ต่ำของการประมวลในรอบนั้น ๆ [28]

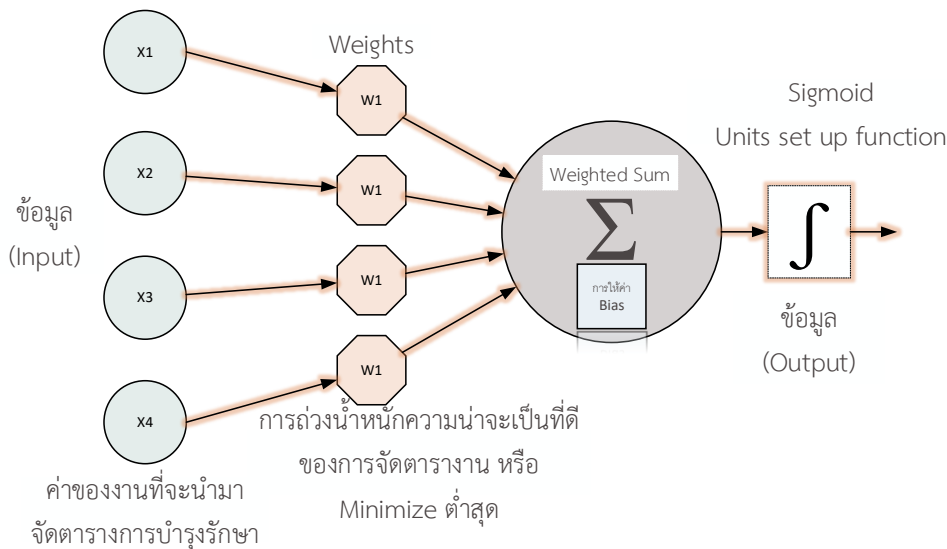
$$\hat{y} = a_j^{(L)} \quad (13)$$

โดยที่ L คือ จำนวนชั้นทั้งหมดในโครงข่าย

3.4.7) การปรับปรุงอย่างต่อเนื่องของการใช้งานโครงข่ายประสาทเทียม ต้องการการปรับปรุงอย่างต่อเนื่อง [28] เนื่องจากข้อมูลใหม่ ๆ จะเข้ามาอย่างสม่ำเสมอ การปรับปรุงโมเดลให้สอดคล้องกับข้อมูลปัจจุบันจะช่วยให้การจัดการการบำรุงรักษาที่มีความแม่นยำและมีประสิทธิภาพมากขึ้น ดังแสดงรูปที่ 1

การใช้โครงข่ายประสาทเทียมมาใช้ในการจัดการการบำรุงรักษา จะช่วยเพิ่มประสิทธิภาพในการวางแผนและบริหารจัดการเวลาของการบำรุงรักษา ซึ่งจะช่วยในการวิเคราะห์ข้อมูล และการคาดการณ์เวลาที่ดียิ่งขึ้นสำหรับการบำรุงรักษา เพื่อลดการหยุดเครื่องจักร ในการผลิต และเพิ่มอายุการใช้งานของเครื่องจักรได้มากขึ้น ซึ่งเทคโนโลยีและอัลกอริทึมที่พัฒนาจะมีความเหมาะสม และนำเข้ามาช่วยในการจัดการกับปัญหาที่มีขนาดใหญ่ และมีความซับซ้อนได้ดี ผู้วิจัยจึงได้พัฒนาโปรแกรมการบำรุงรักษาของบริษัทโดยใช้โปรแกรม VBA ในการพัฒนา ซึ่งโปรแกรม VBA จะมีโครงสร้างการเก็บข้อมูลเครื่องจักร การระบุช่วงเวลาการบำรุงรักษา การระบุช่วงเวลาการเปลี่ยนถ่ายอะไหล่ การจัดการการซ่อมบำรุง และการวิเคราะห์ผลการจัดการการบำรุงรักษากับการวิจัยนี้

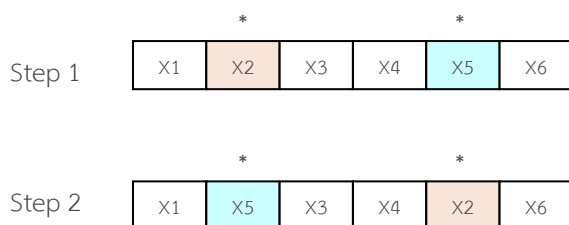
การเก็บข้อมูลแบ่งเป็น 2 ส่วน คือ การเก็บในโปรแกรม และการเก็บเป็นเอกสารการสั่งซ่อมหรือการระบุการซ่อม ซึ่งจะต้องจัดเก็บและเป็นข้อมูลของแผนกบำรุงรักษาของบริษัท



รูปที่ 1 : โครงข่ายประสาทเทียมในการจัดตารางการบำรุงรักษา

3.5) การใช้วิธีการโลคอลเสิร์ช (Local Search: LS)

เป็นเทคนิควิธีการฮิวริสติกที่ใช้ในการหาโซลูชันที่ดีขึ้นโดยการสำรวจโซลูชันปัจจุบัน [29] และเคลื่อนย้ายไปยังโซลูชันใหม่ที่ดีกว่า [20], [21] ในการวิจัยนี้ ได้นำเทคนิคดังกล่าวมาใช้ในการแก้ปัญหาการจัดตารางการบำรุงรักษา [30] โดยเริ่มจากการแปลงปัญหาให้อยู่ในรูปแบบแวนนอนหรือบิท (ตามรูปที่ 2) จากนั้นทำการคำนวณค่าเมคสแปนจากบิทเริ่มต้นเป็นค่าที่ 1 กระบวนการต่อไปจะสุ่มเลือก 2 ตำแหน่งและสลับบิทกันเพื่อสร้างแวนนอนใหม่ แล้วคำนวณค่าเมคสแปนใหม่เป็นค่าที่ 2 หลังจากนั้นนำค่าที่ 1 และค่าที่ 2 มาเปรียบเทียบกับกัน หากค่าเมคสแปนใหม่นั้นน้อยกว่าค่าก่อนหน้า จะเก็บค่าไว้และทำการสุ่มสลับตำแหน่งใหม่ซ้ำไปเรื่อย ๆ จนครบจำนวนที่ตั้งไว้ในโปรแกรม จนได้ค่าเมคสแปนที่ดีที่สุด



รูปที่ 2 : การใช้วิธีการ Local Search ในการแก้ปัญหา

3.6) การหาค่าขอบเขตต่ำสุด (Lower Bound)

เป็นค่าต่ำสุดที่เป็นไปได้สำหรับการจัดลำดับ จะได้ค่าเมคสแปนสำหรับการจัดตารางการบำรุงรักษา หรือค่าใช้จ่ายในปัญหาการ

จัดตาราง การหาค่าใช้วิธีการจัดด้วยวิธีการใด ๆ ในการหาค่าในโปรแกรมที่พัฒนา VBA โดยมีขั้นตอนการหาค่า Lower Bound ที่ทำให้ได้ค่าต่ำสุดและไม่ต่ำกว่านี้

4) ผลการวิจัยและอภิปรายผล

การวิจัยนี้มุ่งเน้นการจัดการจัดการบำรุงรักษาตามตารางการบำรุงรักษา โดยเปรียบเทียบระหว่างระยะเวลา 6 เดือนก่อนการวิจัย (กรกฎาคม-ธันวาคม พ.ศ. 2566) และระยะเวลา 6 เดือนหลังการวิจัย (มกราคม-มิถุนายน พ.ศ. 2567) เพื่อแก้ไขปัญหาการซ่อมบำรุงตามแผนในระบบบำรุงรักษาเชิงป้องกัน (PM) สำหรับเครื่องจักรของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรธานี จากการจัดการบำรุงรักษาต้องดำเนินการกับงานบำรุงรักษาจำนวน n งาน บนหน่วยบำรุงรักษาที่ทำงานอิสระต่อกัน m หน่วย โดยต้องทำงานที่ได้รับมอบหมายให้เสร็จทันเวลาที่กำหนด การจัดลำดับงานให้สอดคล้องกับแผนการบำรุงรักษา PM และเปรียบเทียบกับค่าเมคสแปนต่ำสุดของชุดงาน โดยใช้แบบจำลองกำหนดการเชิงคณิตศาสตร์ นอกจากนี้ยังได้นำโครงข่ายประสาทเทียม มาใช้เพื่อปรับปรุงกระบวนการด้วยอัตราการเรียนรู้ การปรับแต่งพารามิเตอร์ และการร่วมกับวิธีการโลคอลเสิร์ชในการพัฒนาโปรแกรมที่ใช้

การวิเคราะห์ปัญหาที่เกิดขึ้นในแต่ละสัปดาห์ แบ่งออกเป็นปัญหขนาดเล็ก (4×2 , 5×3 , 5×4 , 6×2 , 7×3 , 7×4 , 8×3) และปัญหขนาดใหญ่ (12×5 , 14×2 , 16×3 , 17×3 , 17×4 , 20×5 , 21×5)

โดยการดำเนินการตามแผนการบำรุงรักษาเพื่อลดปัญหาที่พบ และปรับปรุงการทำงานให้มีประสิทธิภาพสูงสุด

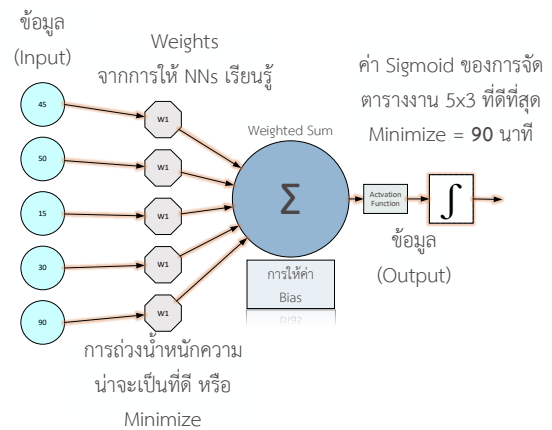
การวิจัยแสดงวิธีการแก้ปัญหางาน 5x3 คือ จำนวนงานที่จะต้องบำรุงรักษาตามแผนจำนวน 5 งาน และมีหน่วยงานที่จะเข้าบำรุงรักษา 3 หน่วยงาน ในการจัดตารางการบำรุงรักษาแสดงตารางที่ 1

ตารางที่ 1 : การแก้ปัญหาการจัดตารางการบำรุงรักษา 5x3

งานที่	ชื่อการบำรุงรักษา	เวลาในการบำรุงรักษา	ต้นทุนในการบำรุงรักษา
5	อัดจาระบีชุดลำเลียง 2	90	780
2	ถ่านน้ำมันหล่อลื่นเครื่องโมแปง 1	50	1,200
1	เปลี่ยนสายพานเครื่องบดแปง	45	960
4	เปลี่ยนเบรคเครื่องหันเส้น	30	569
3	หยอดน้ำมันหล่อลื่นสายพานลำเลียง 1	15	576

จากตารางที่ 1 การแก้ปัญหาการจัดตารางการบำรุงรักษา 5x3 เป็นกระบวนการการแก้ปัญหาด้วย NNs ซึ่งมีการนำข้อมูล (Input) การให้ค่าน้ำหนัก Weights กับ การให้ NNs เรียนรู้ กับ การหาผลรวมของ Weighted Sum กับ การให้ค่า Bias กระบวนการนี้ NNs จะเรียนรู้ค่า และนำไปสู่การหาค่า Sigmoid ของการจัดตารางงาน 5x3 ที่ดีที่สุด Minimize เท่ากับ 90 นาที ที่มีค่าการจัดตาราง คือ 5, 2, 1, 4, 3 มีต้นทุนรวมของการบำรุงรักษานี้เท่ากับ 4,085 บาท ในการจัดตารางการบำรุงรักษานี้

จากปัญหามานำมาเข้ากระบวนการ โครงข่ายประสาทเทียม (แสดงรูปที่ 3) แสดงการใช้ NNs หาค่าคำตอบของการจัดตารางการบำรุงรักษาของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ และการวิจัยได้ออกแบบการทดสอบโปรแกรมกับการหาค่าที่เหมาะสมกับการใช้ NNs การออกแบบการทดลองได้ออกแบบการทดลองเชิงแฟกทอเรียล 2^3 ได้เท่ากับ 8 การทดลอง กับปัญหาขนาดใหญ่ ได้ค่าที่เหมาะสมของค่า Weights ระดับต่ำที่ 0.5 ค่าระดับอัตราการเรียนรู้เท่ากับ 0.01 และการทดลองกับปัญหาขนาดเล็ก ได้ค่าที่เหมาะสมของค่า Weights ระดับต่ำที่ 0.5 ค่าระดับอัตราการเรียนรู้เท่ากับ 0.01 เป็นค่าที่เหมาะสมที่สุดกับการพัฒนาโปรแกรมนี้



รูปที่ 3: โครงข่ายประสาทเทียมกับการจัดตารางงาน

4.1) การจัดตารางการบำรุงรักษาของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ กับปัญหาขนาดเล็ก

โดยการผลิตที่เกิดขึ้นอย่างต่อเนื่องทุกวันทำให้ต้องมีการดำเนินการบำรุงรักษา และไม่ยอมให้มีการ Breakdown หรือถ้ามีต้องน้อยที่สุด การบำรุงรักษาแบบ PM มีการวางแผนการเปลี่ยนอะไหล่ และการจัดตารางการบำรุงรักษาอย่างสม่ำเสมอ ในการวิจัยนี้ได้เก็บข้อมูลเกี่ยวกับปัญหาที่พบในกระบวนการบำรุงรักษา คือ ปัญหาขนาดเล็ก การวิเคราะห์และการเปรียบเทียบข้อมูล ปัญหาขนาดเล็กแสดงไว้ในตารางที่ 2

ตารางที่ 2 : แสดงการเปรียบเทียบผลการวิเคราะห์กับปัญหาขนาดเล็ก

ขนาดของปัญหา (Week)	วิธีการจัดตารางงาน	ค่าขอบเขตต่ำสุด (Lower Bound) (นาที)	ค่าแคสเปน (Makespan) (นาที)	ค่าช่องว่างการเปรียบเทียบ (Gap) (นาที)
4x2	NNs	1,568	1,580	12
	LS		1,578	10
5x3	NNs	1,740	1,740	0
	LS		1,802	62
5x4	NNs	2,102	2,102	0
	LS		2,102	0
6x2	NNs	2,089	2,089	0
	LS		2,102	13
7x3	NNs	2,358	2,358	0
	LS		2,358	0
7x4	NNs	2,689	2,689	0
	LS		2,689	0
8x3	NNs	3,329	3,329	0
	LS		3,335	6

จากตารางที่ 2 แสดงการเปรียบเทียบผลการวิเคราะห์กับปัญหาขนาดเล็ก พบว่า การใช้วิธีโครงข่ายประสาทเทียมและวิธีการโลคอลเสิร์ช สามารถหาคำตอบในการจัดตารางการบำรุงรักษาของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ ได้อย่างมีประสิทธิภาพ ตัวอย่างเช่น ปัญหาขนาด 7x3 ซึ่งให้ค่าแมคสแปนเท่ากับ 2,358 นาที เมื่อเปรียบเทียบกับค่า Lower Bound ที่เท่ากับ 2,358 นาที จะเห็นว่าค่าช่องว่างการเปรียบเทียบ (Gap) เป็น 0 นั้นหมายความว่า การจัดการกับปัญหาขนาดเล็กด้วย NNs และ LS ให้ผลลัพธ์ที่ดี

4.2) การจัดตารางการบำรุงรักษาของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ กับปัญหาขนาดใหญ่

ในกระบวนการบำรุงรักษาที่วางแผนร่วมกับการจัดตารางและนำแบบจำลองกำหนดการเชิงคณิตศาสตร์มาใช้ร่วมกับการใช้ NNs และ LS การเปรียบเทียบข้อมูลปัญหาขนาดใหญ่ดังกล่าว แสดงไว้ในตารางที่ 3

ตารางที่ 3 : แสดงการเปรียบเทียบผลการวิจัยกับปัญหาขนาดใหญ่

ขนาดของปัญหา (Week)	วิธีการการจัดตารางงาน	ค่าขอบเขตต่ำสุด (Lower Bound) (นาที)	ค่าแมคสแปน (Makespan) (นาที)	ค่าช่องว่างการเปรียบเทียบ (Gap) (นาที)
15x5	NNs	1,950	1,950	0
	LS		1,950	0
14x2	NNs	2,240	2,242	2
	LS		2,244	4
16x3	NNs	2,309	2,327	18
	LS		2,337	28
17x3	NNs	2,540	2,540	0
	LS		2,582	42
17x4	NNs	2,785	2,785	0
	LS		2,759	26
20x4	NNs	3,689	3,689	0
	LS		3,689	0
21x4	NNs	3,890	3,890	0
	LS		3,905	15

จากตารางที่ 3 แสดงการเปรียบเทียบผลการวิจัยกับปัญหาขนาดใหญ่ พบว่า การใช้วิธีโครงข่ายประสาทเทียม NNs และ LS สามารถหาคำตอบในการจัดตารางการบำรุงรักษาของบริษัท

ผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์ สำหรับปัญหาขนาดใหญ่ได้อย่างมีประสิทธิภาพ ตัวอย่างเช่น ปัญหาขนาด 20x4 ซึ่งให้ค่าแมคสแปนเท่ากับ 3,689 นาที เมื่อเปรียบเทียบกับค่า Lower Bound ที่เท่ากับ 3,689 นาที จะเห็นว่าค่าช่องว่างการเปรียบเทียบ (Gap) เป็น 0 นั้นหมายความว่า การจัดการกับปัญหาขนาดใหญ่ด้วย NNs และ LS ให้ผลลัพธ์ที่ดีเช่นกัน

4.3) ผลการแก้ปัญหาการบำรุงรักษาของบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตถ์

4.3.1) การวิจัยได้ทำการเก็บและการวิเคราะห์ข้อมูลก่อนการวิจัย โดยมีการจัดเก็บข้อมูลล่วงหน้าเป็นระยะเวลา 6 เดือน ข้อมูลก่อนการวิจัยถูกเก็บระหว่างเดือนกรกฎาคม ถึงเดือนธันวาคม 2566 ผลการวิเคราะห์แสดงในตารางที่ 4

ตารางที่ 4: ข้อมูลการบำรุงรักษาก่อนดำเนินการวิจัยของโรงงาน กรณีศึกษา

เดือน	จำนวนงานแก้ปัญหาบำรุงรักษา (งาน)	ระยะเวลาในการบำรุงรักษา (นาที)	ต้นทุนการแก้ปัญหา (บาท)	ร้อยละการบำรุงรักษาต่อระยะเวลาในการบำรุงรักษา
ก.ค. 2566	86	8,964	282,265	13.38
ส.ค. 2566	127	9,875	254,684	14.74
ก.ย. 2566	168	10,821	345,669	16.16
ต.ค. 2566	159	12,412	451,785	18.53
พ.ย. 2566	184	11,458	545,897	17.11
ธ.ค. 2566	171	13,427	457,158	20.05
รวม		66,957	2,337,458	

จากตารางที่ 4 ข้อมูลการบำรุงรักษาก่อนดำเนินการวิจัยของโรงงาน กรณีศึกษา พบว่า มีจำนวนระยะเวลาในการบำรุงรักษา รวมเท่ากับ 66,957 นาที มีต้นทุนในการบำรุงรักษาจำนวน 6 เดือน รวมเท่ากับ 2,337,458 บาท

4.3.2) การวิจัยได้ทำการเก็บและการวิเคราะห์ข้อมูลหลังการวิจัย การจัดเก็บข้อมูลหลังดำเนินการ โดยมีการจัดเก็บข้อมูลเป็นระยะเวลา 6 เดือน ข้อมูลหลังการวิจัยถูกเก็บระหว่างเดือนมกราคม ถึงเดือนมิถุนายน 2567 ดังแสดงตารางที่ 5

ตารางที่ 5: ข้อมูลการบำรุงรักษาหลังดำเนินการวิจัยของโรงงาน กรณีศึกษา

เดือน	จำนวนงานแก้ปัญหา บำรุงรักษา (งาน)	ระยะเวลาในการ บำรุงรักษา (นาท)	ต้นทุนการแก้ปัญหา (บาท)	ร้อยละการบำรุงรักษาต่อ ระยะเวลาในการ บำรุงรักษา
ม.ค. 2567	145	7,642	241,258	15.10
ก.พ. 2567	141	7,423	241,752	14.66
มี.ค. 2567	159	6,125	255,641	12.10
เม.ย. 2567	164	9,874	254,147	19.51
พ.ค. 2567	191	10,954	358,144	21.64
มิ.ย. 2567	179	8,584	564,120	16.96
รวม		50,602	1,915,062	

จากตารางที่ 5 ข้อมูลการบำรุงรักษาหลังดำเนินการวิจัยของโรงงาน กรณีศึกษา พบว่า มีจำนวนระยะเวลาในการบำรุงรักษา รวมเท่ากับ 50,602 นาที่ มีต้นทุนในการบำรุงรักษาจำนวน 6 เดือน รวมเท่ากับ 1,915,062 บาท

4.3.3) การเปรียบเทียบผลการวิเคราะห์ การวิจัยได้ทำการเปรียบเทียบผลลัพธ์ก่อนและหลังการดำเนินการ พบว่า ก่อนการวิจัยมีระยะเวลาในการบำรุงรักษารวมทั้งสิ้น 66,957 นาที่ หลังการวิจัยมีระยะเวลาในการบำรุงรักษาลดลงเหลือ 50,602 นาที่ ซึ่งลดลง 16,355 นาที่ หรือคิดเป็นร้อยละ 13.91 และใน ด้านต้นทุนการบำรุงรักษาลดลง 6 เดือน ก่อนการวิจัยมีต้นทุนรวมในการบำรุงรักษาอยู่ที่ 2,337,458 บาท หลังการวิจัยมี ต้นทุนรวมลดลงเหลือ 1,915,062 บาท ลดลง 422,396 บาท หรือคิดเป็นร้อยละ 9.93 การเปรียบเทียบนี้แสดงให้เห็นถึงการ ลดลงอย่างมีนัยสำคัญทั้งในด้านระยะเวลาและต้นทุนการ บำรุงรักษาหลังจากการดำเนินการวิจัย

4.4) การเปรียบเทียบผลค่าระยะเวลาเฉลี่ยระหว่างการเกิด เหตุขัดข้องของเครื่องจักร (MTBF)

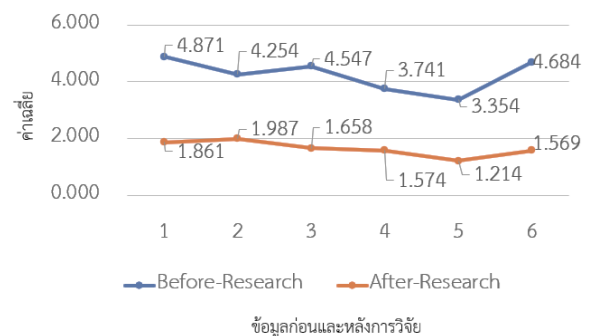
การวิจัยได้เก็บข้อมูลก่อนการวิจัย 6 เดือน และหลังการวิจัย 6 เดือน พบว่า ค่า MTBF ก่อนการวิจัย มีค่าเท่ากับ 48.57 ชั่วโมง ขณะที่หลังการวิจัยค่า MTBF เพิ่มขึ้นเป็น 83.84 ชั่วโมง ซึ่ง แสดงให้เห็นว่าค่า MTBF เพิ่มขึ้นร้อยละ 35.27 เมื่อเปรียบเทียบ ระหว่างก่อนและหลังการวิจัย

4.5) การเปรียบเทียบค่าเวลาเฉลี่ยในการซ่อมแซม (Mean Time to Repair: MTTR)

จากข้อมูลวิธีการแก้ปัญหาที่พัฒนาขึ้น ทำการเก็บผลการ วิเคราะห์ก่อนการวิจัย 6 เดือน และหลังการวิจัย 6 เดือนแสดง ให้เห็นว่า ค่าเฉลี่ยระหว่างการเกิดเหตุขัดข้องก่อนการวิจัยอยู่ที่ 3.49 ชั่วโมง ขณะที่ค่าเฉลี่ยหลังการวิจัยลดลงเหลือ 1.36 ชั่วโมง ซึ่งสะท้อนถึงการลดลงของค่า MTTR ในแต่ละเดือนหลังการวิจัย คิดเป็นร้อยละ 30.48

4.6) การเปรียบเทียบค่าร้อยละของเวลาที่เครื่องจักรเกิดเหตุ ขัดข้อง (Average Machine Downtime)

การวิจัยได้เก็บข้อมูลเกี่ยวกับเวลาที่เครื่องจักรเกิดเหตุขัดข้อง ระหว่างการผลิต โดยบันทึกจำนวนครั้งที่เครื่องจักรเสียหายก่อน การวิจัย 6 เดือน และหลังการวิจัย 6 เดือน จากนั้นทำการ คำนวณค่าเวลาเป็นค่าร้อยละ (สมการที่ 3) เมื่อเทียบกับเวลาใน การผลิตทั้งหมดของผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในแต่ละเดือน ผลการวิเคราะห์แสดงในรูปที่ 4 ซึ่งเปรียบเทียบค่าร้อยละของ เวลาที่เครื่องจักรขัดข้องก่อนและหลังการวิจัย พบว่า ค่าร้อยละ ของเวลาที่ใช้ในการบำรุงรักษาในแต่ละเดือนเมื่อเทียบกับเวลาใน การบำรุงรักษาทั้งหมดลดลงหลังการวิจัย โดยคิดเป็นร้อยละ 44.14



รูปที่ 4 : ค่าร้อยละของเวลาที่เครื่องจักรเกิดเหตุขัดข้อง

5) สรุปผล

การวิจัยนี้มุ่งเน้นการจัดตารางการบำรุงรักษาของบริษัทผลิต และจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรดิตต์ โดยได้จัดเก็บข้อมูลปัญหาขนาดเล็ก และขนาดใหญ่ ในการบำรุง รักษาแบบ PM เพื่อป้องกันการ Breakdown หรือให้เกิดขึ้นน้อย ที่สุด ข้อมูลถูกเก็บระหว่างเดือนกรกฎาคม-ธันวาคม พ.ศ. 2566 (ก่อนการวิจัย) และเดือนมกราคม-มิถุนายน พ.ศ. 2567 (หลัง

การวิจัย) สำหรับปัญหาขนาดเล็ก เช่น 7x3 ซึ่งให้ค่าเมคสแปนเท่ากับ 2,358 นาที เมื่อเปรียบเทียบกับค่า Lower Bound ที่เท่ากับ 2,358 นาที ทำให้ Gap เป็น 0 สำหรับปัญหาขนาดใหญ่ เช่น 20x4 ซึ่งให้ค่าเมคสแปน เท่ากับ 3,689 นาที เมื่อเปรียบเทียบกับค่า Lower Bound ที่เท่ากับ 3,689 นาที จะเห็นว่าค่าช่องว่างการเปรียบเทียบเป็น 0 แสดงถึงประสิทธิภาพในการแก้ปัญหาค่าใหญ่ด้วย NNs และ LS ในปัญหาทั้งขนาดเล็กและขนาดใหญ่ หลังการวิจัย พบว่า ระยะเวลาการบำรุงรักษารวมลดลงจาก 66,957 นาที เป็น 50,602 นาที และต้นทุนการบำรุงรักษาลดลงจาก 2,337,458 บาท เป็น 1,915,062 บาท ลดลงจากเดิมเป็นจำนวนเงิน 422,396 บาท หรือคิดเป็นร้อยละ 9.93 และการวิจัยได้เปรียบเทียบผลค่าระยะเวลาเฉลี่ยระหว่างการเกิดเหตุขัดข้องของเครื่องจักร พบว่าค่า MTBF เพิ่มขึ้นจาก 48.57 ชั่วโมงก่อนการวิจัยเป็น 83.84 ชั่วโมง หลังการวิจัยเพิ่มขึ้นร้อยละ 35.27 ส่วนค่า MTTR ลดลงจาก 3.49 ชั่วโมงเป็น 1.36 ชั่วโมง ลดลงร้อยละ 30.48 นอกจากนี้ ค่าเวลาที่เครื่องจักรเกิดเหตุขัดข้องลดลง ร้อยละ 44.14 หลังการวิจัย

6) ข้อเสนอแนะ

ผลการวิจัยนี้สามารถใช้เป็นแนวทางในการจัดตารางการบำรุงรักษาเพื่อใช้เพิ่มประสิทธิภาพในกระบวนการผลิต เสนอให้มีการฝึกอบรมบุคลากรอย่างต่อเนื่องในด้านการใช้เทคโนโลยีใหม่ ๆ ที่นำมาใช้ในโรงงานร่วมกับการวิเคราะห์ข้อมูล และเพิ่มอัตราการเรียนรู้ให้กับโปรแกรมที่พัฒนาอย่างต่อเนื่อง

กิตติกรรมประกาศ

ขอบคุณบริษัทผลิตและจำหน่ายผลิตภัณฑ์อาหารกึ่งสำเร็จรูปในจังหวัดอุดรธานี และหลักสูตรเทคโนโลยีอุตสาหกรรมที่สนับสนุนการวิจัย การใช้พื้นที่วิจัย เครื่องมือ เครื่องจักร ซอฟต์แวร์ และอุปกรณ์ในการวิจัยในครั้งนี้

REFERENCES

- [1] A. N. Mustapha, Y. Zhang, Z. Zhang, Y. Ding, Q. Yuan and Y. Li (2021). "Taguchi and ANOVA analysis for the Optimization of the microencapsulation of a volatile phase change material," *J. Mat. Res. Technol.*, vol. 11, pp. 667–680, 2021.
- [2] R. H. Myers and D. C. Montgomery, *Response Surface Methodology: Process and Product Optimization Using Designed Experiments*. New York, NY, USA: John Wiley & Sons, 1995.
- [3] H.-K. Hwang and S.-J. Kim, "Investigation on the effective factor calculation of electropolishing using full factorial design and mechanism model by microscopic analysis for super austenitic stainless steel," *Surfaces and Interfaces*, vol. 37, 2023, Art. no. 102730.
- [4] D. Grewal, S. Benoit, S. M. Noble, A. Guha, C.-P. Ahlbom, and J. Nordfält, "Leveraging in-store technology and AI: Increasing customer and employee efficiency and enhancing their experiences," *J. Retail.*, vol. 99, no. 4, pp. 487–504, 2023.
- [5] A. Jankovic, G. Chaudhary, and F. Goia, "Designing the design of experiments (DOE) – An investigation on the influence of different factorial designs on the characterization of complex systems," *Energy Build.*, vol. 250, Nov. 2021, Art. no. 111298.
- [6] H. Wu, H. Chen, L. Fan, P. Pan, G. Xu, and L. Wu, "Performance analysis of a novel co-generation system integrating a small modular reactor and multiple hydrogen production equipment considering peak shaving," *Energy*, vol. 302, Sep. 2024, Art. no. 131887.
- [7] S. Tang, J. Ma, Z. Yan, Y. Zhu, and B. C. Khoo, "Deep transfer learning strategy in intelligent fault diagnosis of rotating machinery," *Eng. Appl. Artif. Intell.*, vol. 134, Aug. 2024, Art. no. 108678.
- [8] H. Corrotea, H. Portales, L. Amigo, G. Gatica, A. T. Palacio, D. Mondragon, and M. Ramos, "Maintenance process analysis in a port cargo company through discrete event simulation," *Procedia Comput. Sci.*, vol. 231, pp. 415–420, 2024.
- [9] P. Mukherjee, R. Das, and M. Jawed, "Assessment of routine operation and maintenance work of a small community water treatment plant in North Guwahati, Assam, India," *Groundwater for Sustain. Develop.*, vol. 25, May 2024, Art. no. 101123.
- [10] U. Rai, G. Oluleye, and A. Hawkes, "Stochastic optimisation model to determine the optimal contractual capacity of a distributed energy resource offered in a balancing services contract to maximise profit," *Energy Rep.*, vol. 11, pp. 5800–5818, Jun. 2024.

- [11] L. Cheng, Q. Tang, and L. Zhang, "Production costs and total completion time minimization for three-stage mixed-model assembly job shop scheduling with lot streaming and batch transfer," *Eng. Appl. Artif. Intell.*, vol. 130, Apr. 2024, Art. no. 107729.
- [12] Z. Rui *et al.* "Graph reinforcement learning for flexible job shop scheduling under industrial demand response: A production and energy nexus perspective," *Comput. Ind. Eng.*, vol. 193, Jul. 2024, Art. no. 110325.
- [13] B. Du, S. Han, J. Guo, and Y. Li, "A hybrid estimation of distribution algorithm for solving assembly flexible job shop scheduling in a distributed environment," *Eng. Appl. Artif. Intell.*, vol. 133, Jul. 2024, Art. no. 108491.
- [14] H. Shuihua, Z. Ling, C. Kui, Z. Luo, and D. Mishra, "Appointment scheduling and routing optimization of attended home delivery system with random customer behavior," *Eur. J. Oper. Res.*, vol. 262, no. 3, pp. 966–980, Nov. 2017.
- [15] A. Choo, Y. Xia, G. P. Zhang, and C. Liao, "Leader behavioral integrity for safety and its impact on worker preventive maintenance behavior and operational performance," *Saf. Sci.*, vol. 177, Sep. 2024, Art. no. 106577.
- [16] M. D. O. Dos Reis, R. Godina, C. Pimented, F. J. G. Silva, and J. C. O. Matiasac, "A TPM strategy implementation in an automotive production line through loss reduction," *Procedia Manuf.*, vol. 38, pp. 908–915, 2019.
- [17] W. Zhang, J. Gan, S. He, T. Li, and Z. He, "An integrated framework of preventive maintenance and task scheduling for repairable multi-unit systems," *Rel. Eng. Syst. Saf.*, vol. 247, Jul. 2024, Art. no. 110129.
- [18] J. L. M. de Andrade, M. J. F. Souza, E. M. de Sa, G. C. Menezes, and S. R. de Souza, "Formulations and heuristic for the long-term preventive maintenance order scheduling problem," *Comput. Oper. Res.*, vol. 170, Oct. 2024, Art. no. 106781.
- [19] T. Yuan and X. Yan, "Application analysis of heuristic algorithms integrating dynamic programming in RNA secondary structure prediction," *Intell. Syst. Appl.*, vol. 23, Sep. 2024, Art. no. 200400.
- [20] A. Phuk-in, "Production-scheduling problem: A case of SD Tractors Company, Limited," in *Proc. 40th Conf. Ind. Eng. Netw. (IE Network)*, Nakhon Pathom, Thailand, May 2022, pp. 469–474.
- [21] A. Phuk-in, "Production planning and machine maintenance schedule of Dragon Green Energy Company, Limited," *J. Ind. Technol.*, vol. 20, no. 1, pp. 62–80, 2024.
- [22] C. Muangdit and K. Lurang, "Optimized PIDA controller design with bat-inspired algorithm for temperature control of electric furnace system," (in Thai), *J. Eng. Digit. Technol. (JEDT)*, vol. 12, no. 1, pp. 67–74, 2024.
- [23] P. Alavian, Y. Eun, K. Liu, S. M. Meerkov, and L. Zhang, "The (α, β) -precise estimates of MTBF and MTTR: definitions, calculations, and induced effect on machine efficiency evaluation," *IFAC-PapersOnLine*, vol. 52, no. 13, pp. 1004–1009, 2019.
- [24] A. Khamaj, A. M. Ali, R. Saminathan, and M. Shanmugasundaram, "Human factors engineering simulated analysis in administrative, operational and maintenance loops of nuclear reactor control unit using artificial intelligence and machine learning techniques," *Heliyon*, vol. 10, no. 10, May 2024, Art. no. e30866.
- [25] R. Zheng, M. Liu, Y. Zhang, Y. Wang, and T. Zhong, "An optimization method based on improved ant colony algorithm for complex product change propagation path," *Intell. Syst. Appl.*, vol. 23, Sep. 2024, Art. no. 200412.
- [26] J. Zhao, H. Mao, P. Mao, and J. Hao, "Learning path planning methods based on learning path variability and ant colony optimization," *Syst. Soft Comput.*, vol. 6, Dec. 2024, Art. no. 200091.
- [27] J. Bai, G.-R. Liu, T. Rabczuk, Y. Wang, X.-Q. Feng, and Y. T. Gu, "A robust radial point interpolation method empowered with neural network solvers (RPIM-NNs) for nonlinear solid mechanics," *Comput. Methods Appl. Mechanics Eng.*, vol. 429, Sep. 2024, Art. no. 1171591.
- [28] Z. Liang, D. Ren, B. Xue, J. Wang, W. Yang, and W. Liu, "Verifying safety of neural networks from topological perspectives," *Sci. Comput. Program.*, vol. 236, Sep. 2024, Art. no. 103121.
- [29] Y. Xu, X. Li, and X. Meng, "A Q-learning based iterated local search algorithm for human-UAV cooperation in restoring transmission network," *Expert Syst. Appl.*, vol. 252, Oct. 2024, Art. no. 124200.
- [30] A. E. Yahiaoui, S. Afifi, and H. Allaoui, "Enhanced iterated local search for the technician routing and scheduling problem," *Comput. Oper. Res.*, vol. 160, Dec. 2023, Art. no. 106385.



Transforming Unstructured Data in IT Project: A Comparative Study of Zero-Shot and Generative AI Text Classification

Cai Tung-lersloy^{1*} Worapat Paireekreng² Nantika Prinyapol³

^{1*,2,3}*College of Innovative Technology and Engineering, Dhurakij Pundit University, Bangkok, Thailand*

*Corresponding Author. E-mail address: 637191110006@dpu.ac.th

Received: 7 August 2024; Revised: 7 October 2024; Accepted: 8 October 2024

Published online: 26 December 2025

Abstract

In today's world, we have a lot of messy, unorganized data from things like comments, interviews, and images. This is especially true in IT projects, where there's often too much information to handle easily. Our study looks at how we can turn this messy data into useful numbers and insights using smart computer programs. We tested two main methods: Zero-Shot Text Classification and Generative AI Text Classification. Zero-Shot is like having a smart assistant that can sort information without needing examples first. Generative AI is more like having a creative writer who can come up with new examples to help sort information. We asked 42 participants with experience in working with unstructured data to answer some questions, then used these methods to analyze their answers. We found that Zero-Shot works better for information that has clear patterns, while Generative AI is good at handling more complex or unclear information. Our results show that choosing the right method can make a big difference in how well we understand and use the data. Zero-Shot was about 15% more accurate for well-organized information, while Generative AI was 20% better at dealing with complex, messy data. This research helps companies and researchers choose the best way to make sense of their data, especially in IT projects where there's often too much information to handle manually.

Keywords: Qualitative data, Unstructured data, Zero-shot text classification

I. INTRODUCTION

The proliferation of unstructured qualitative data in the digital era poses significant challenges for effective analysis, particularly in information technology (IT) projects. With approximately 80% of data remaining unstructured [1], this issue extends across various sectors. IT projects, encompassing software development, network upgrades, and cybersecurity implementations, are at the forefront of managing this data deluge.

Transforming unstructured data into quantitative insights involves unitization, categorization, and coding. Large Language Models (LLMs) like GPT-3 have shown promise in automating this process [2], though careful application is crucial to avoid inaccuracies. Zero-Shot Classification offers powerful tools for categorizing data without task-specific training. The CLORE (Classification by LOGical REasoning) framework [3] and semantic knowledge integration techniques [4] exemplify this approach, leveraging logical reasoning on natural language explanations for effective classification.

This paper explores the application of LLMs and Zero-Shot Classification in revolutionizing data management for unstructured data in IT projects. By leveraging AI to transform qualitative data into quantitative insights, organizations can enhance decision-making and data analysis efficiency [5], unlocking the economic and innovative potential of unstructured data.

II. LITERATURE REVIEW

This review examines Zero-Shot Text Classification and Generative AI Text Classification as key methodologies for transforming unstructured data in IT projects. IT projects, in this context, refer to technology-driven initiatives within organizations involving the development, implementation, or maintenance of information systems and digital infrastructure. These encompass activities such as software development, network upgrades, data

management systems, cybersecurity implementations, cloud migrations, ERP and CRM system deployments, and mobile application development. Such projects often generate and handle vast amounts of unstructured data, making them ideal candidates for advanced AI-driven analysis techniques. Zero-Shot Classification utilizes existing knowledge to categorize text without task-specific training [3], [4], while Generative AI Text Classification employs large language models to generate and classify text, adapting to complex patterns [2]. Recent advancements by Zhang et al. [4] and Abburi et al. [2] have enhanced these techniques, with Ye et al. [6] expanding Zero-Shot capabilities using pre-trained models and prompt learning. Despite the potential demonstrated by Yin et al. [7] and Brown et al. [8] in NLP tasks and human-like text generation, the comparative effectiveness of these techniques in transforming qualitative IT project data into quantitative insights remains unexplored. This study aims to bridge this research gap, potentially revolutionizing unstructured data processing and analysis in IT project management.

A. Transforming Qualitative Data into Quantitative Results

The process of converting qualitative data into quantitative insights is crucial for effective analysis in various fields. Srnka and Koeszegi [9] propose a systematic approach involving structured data collection, rigorous transcription, and well-defined categorization and coding processes. This method emphasizes the scientific measurement of qualitative data, although it acknowledges that natural human understanding and open-ended responses often yield powerful qualitative insights. Modern AI techniques, including Zero-Shot Text Classification [3], Large Language Models (LLMs) [2], and Generative AI Text Classification [4], offer promising solutions to automate and enhance this transformation process, addressing the challenges

posed by large datasets and the need for efficient data analysis.

As Table 1 illustrates, there are various mixed research designs for integrating qualitative and quantitative approaches:

Table 1: Qualitative-Quantitative Research Designs: Types, Descriptions, and Aims

Research	Description
Sequential two-studies design	Description: Qualitative data and quantitative data are collected and analyzed in sequential order. Aim: Investigate under-researched fields, develop hypotheses or create instruments for subsequent quantitative measurement, or provide explanations.
Concurrent two-studies design	Description: Both quantitative and qualitative data are collected and analyzed in separate procedures. Aim: Cross-validate or corroborate findings of the two approaches.
Integrated elaboration design	Description: Quantitative data is analyzed using qualitative procedures. Aim: Investigate and understand the problem in depth, derive new theoretical insights.
Integrated generalization design	Description: Qualitative material is collected and transformed into categorical data for further quantitative analysis. Aim: Derive both theory and generalizable results.

B. Zero-Shot Text Classification

Zero-Shot Text Classification is an advanced technique for categorizing text without task-specific training. The CLORE (Classification by LOGical REasoning) framework, introduced by Han et al. [3], exemplifies this approach through two main stages:

1) Logical Parsing: Breaking down explanations into logical structures to identify relevant attributes.

2) Logical Reasoning: Matching these attributes to input data for classification scoring.

This method demonstrates superior performance in tasks requiring high-level logical reasoning and offers improved interpretability compared to baseline models. It also shows robustness against linguistic biases, making it versatile for various classification [3]. Zhang et al. [4] further enhanced this approach by proposing a two-phase framework that integrates semantic knowledge:

1) Coarse-Grained Classification: Using a traditional classifier to determine if an input belongs to seen or unseen classes.

2) Fine-Grained Classification: Employing a zero-shot classifier with semantic knowledge-based feature augmentation for more precise categorization.

This framework significantly improves zero-shot text classification, particularly for domains with evolving or diverse classification needs. However, it may face challenges with tasks requiring new data generation or handling highly complex patterns, highlighting the potential complementary nature of zero-shot and generative approaches in addressing diverse text classification challenges.

C. Generative AI Text Classification

Generative AI text classification differs from zero-shot text classification by focusing on the generation and classification of new text data, rather than relying solely on existing data and logical reasoning. Generative AI models, such as those combining multiple LLMs, can create new content and classify it based on learned patterns. [2] explore the application of ensemble models combining multiple Large Language Models (LLMs) for generative AI text classification.

Table 2 from shows the results of various models for the Binary-English task. The ensemble with voting

classifier outperformed individual models, demonstrating the strength of combining outputs from multiple LLMs.

Table 2: Results for the Binary-English task

Model	Accuracy	Macro F1	Precision	Recall
deberta-large				
	0.62	0.546	0.783	0.61
xlm-r-100langs-bert-base-nli-stsb-mean-tokens				
	0.647	0.592	0.782	0.639
roberta-base-openai-detector				
	0.679	0.636	0.805	0.671
xlm-roberta-large-xnli-anli				
	0.618	0.543	0.782	0.608
roberta-large				
	0.623	0.551	0.784	0.613
Ensemble with Voting classifier				
	0.751	0.733	0.826	0.745

Generative AI text classification excels in handling complex and nuanced tasks by leveraging LLMs' capabilities to create new data, filling gaps in existing datasets. However, this approach requires careful management to avoid generating inaccurate or misleading information. Recent research by Abi Akl [10] demonstrates the synergy between traditional ML techniques and LLM-generated data, achieving a Macro-F1 score of 88.401% on seed data with a 1.5% performance increase when augmented by LLM-generated data. This combination not only enhances text classification accuracy but also provides a robust pathway for extracting and utilizing unstructured data within larger frameworks. The effectiveness of LLMs in creating robust text embeddings further expands their application, such as in code comment classification. Models like ChatGPT showcase the potential of generating additional data to improve classical machine learning systems, highlighting the versatility of AI in handling diverse unstructured data types.

D. The Framework of Extracting Unstructured Usage for Big Data Platform

To test AI techniques in transforming qualitative data, [11] propose a framework for extracting and utilizing unstructured data in organizations, distinguishing between structured and unstructured data. Chasupa and Paireekreng's framework [11] enhances decision-making by converting qualitative inquiries into quantitative measurements through a 4x4 questionnaire, ensuring effective data transformation.

Table 3: Unstructured Big Data Extracting Model

		Unstructured Data Form			
		Object	Event	Command	Outcome
Unstructured Data Activity	Property				
	Cause				
	Truth				
	Journey				

Literature Review Summary: The discussed research highlights that for transforming qualitative data into quantitative results, the framework provided by Srnka and is effective for smaller datasets [9]. However, when dealing with large datasets, retrospective interviews, or comments, significant challenges and limitations arise, especially in research grounded in qualitative data or in extracting data from new domains. Replacing the three stages of Srnka and Koeszegi's framework with AI techniques can enhance the diversity and scope of analysis and research.

Zero-Shot Text Classification and Generative AI Text Classification offer promising solutions. Zhang et al. show that Zero-Shot Text Classification is highly effective for tasks requiring logical parsing and semantic knowledge integration. This method is suitable for structured, interconnected data. On the other hand, Generative AI Text Classification, as explored by Abburi et al. [2] and Abi Akl [10], is adept at handling complex and nuanced tasks by generating new data and enhancing traditional ML systems with LLM-generated data. This method is beneficial for data requiring interpretation or filling in gaps.

To determine the most appropriate method and validate the research hypothesis, both tools will be employed to compare their results. The goal is to ascertain which method is more accurate and suitable for converting large-scale qualitative data into quantitative insights and to identify the specific contexts in which each method excels. This approach will provide a comprehensive understanding of the effectiveness of AI in managing and transforming qualitative data.

III. RESEARCH METHODOLOGY

The proposed methodology aims to transform qualitative data into quantitative insights using two AI techniques: Zero-Shot Text Classification and Generative AI Text Classification. The methodology will involve the following key steps as Figure 1.

- 1) Data Collection: Gather qualitative data through detailed questionnaires and narrative responses.
- 2) Data Wrangling: Clean and transform the collected data to ensure consistency and usability.
- 3) Data Preprocessing: Prepare the qualitative data for analysis by categorizing and coding it.
- 4) Chose AI Technique is Model Application.

5) Zero-Shot Text Classification: Apply Zero-Shot Text Classification to categorize the qualitative data without task-specific training.

6) Generative AI Text Classification: Apply Generative AI Text Classification using ensemble models combining multiple LLMs to generate new data and classify it.

7) Comparison and Analysis: Compare the results of both AI techniques to determine their accuracy, effectiveness, and suitability for various data contexts.

8) Evaluation: Evaluate the performance of each AI technique and identify the specific contexts in which each method excels.

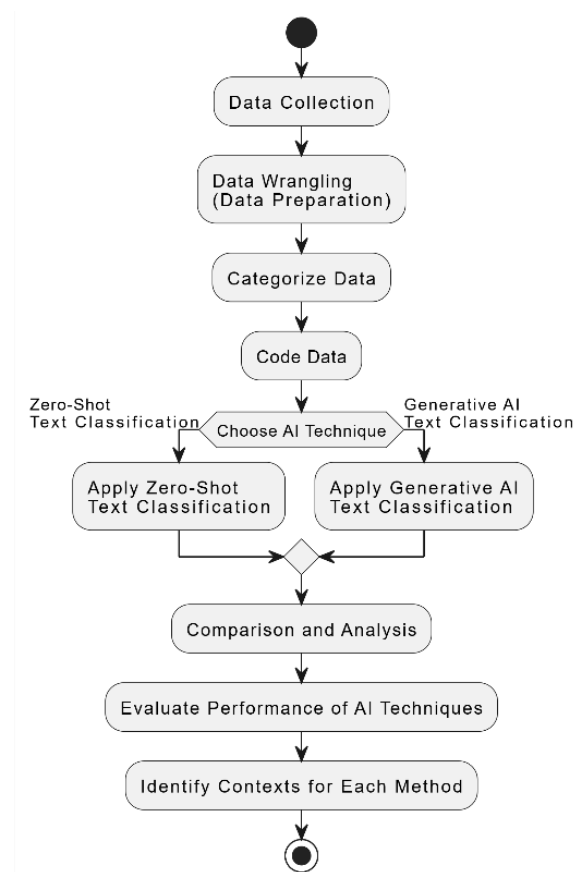


Figure 1: Research Methodology

A. Data Collection

The data collection process will involve gathering qualitative data from participants involved in projects related to unstructured data. The participants will be selected from nine different types of projects and six IT

job positions from Talence [12]. The total population size will be 54 individuals, and the sample size will be calculated using Cochran's formula to ensure representativeness.

B. Data Wrangling.

Data wrangling is vital for preparing qualitative data for analysis, involving tasks like handling missing values, correcting inconsistencies, and segmenting narrative responses. Language models, particularly large ones, assist in these tasks through few-shot or zero-shot inference [13]. Participants are drawn from projects in various domains, including 1) OCR projects, 2) NLP projects, 3) Paperless document management systems, 4) Sentiment and opinion analysis, 5) Image and video analysis, 6) Audio data analysis, 7) Social media data analysis, and 8) Recommendation systems. The appropriate sample size, determined using Cochran's formula by bin Ahmad and binti Halim [14], with an 85% confidence level, is 42 participants as shown in Equation 1.

$$n = \frac{207.36}{1 + \left(\frac{207.36-1}{54}\right)} = \frac{207.36}{1+3.83} = \frac{207.36}{4.83} \approx 42.95 \quad (1)$$

IV. RESULTS

This section presents the findings of the study, focusing on the transformation of qualitative data into quantitative insights using Zero-Shot Text Classification and Generative AI Text Classification. The objectives of the research are reiterated to provide context for the results. A detailed analysis of each table is presented, followed by an integrated discussion of all results to provide a comprehensive overview of the study's findings.

A. Data Preparation and Initial Observations.

To ensure clarity and diversity in the questions, the 4x4 model from "The Framework of Extracting Unstructured Usage for Big Data Platform" was extended

to include dimensions of project management. This addition enriched the data with project management perspectives, resulting in a comprehensive 4x4x4 framework. The added dimensions are.

- 1) Time: Represents different phases of a project, helping to identify when events or activities occur and estimate the duration required for each phase.
- 2) Priority: Indicates the importance level of data or activities, aiding in prioritizing urgency and resource allocation efficiently.
- 3) Resources: Represents the resources required for operations or analysis, such as data, technology, personnel, or capital.
- 4) Stakeholder: Identifies groups affected by or influential to the project, helping to manage expectations and requirements clearly [1].

This created a set of questions with enhanced depth, forming a 4x4x4 framework as illustrated in Figure 1.

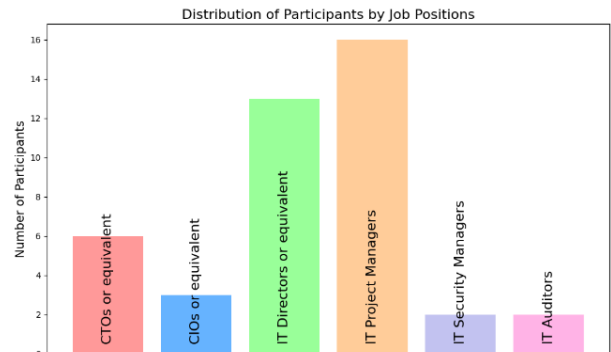


Figure 2: Distribution of Participants by Job Positions

The study involved participants from nine types of unstructured data projects and six IT job positions, totaling 54 individuals. Using Cochran's formula, the sample size was determined to be 42-43 participants, resulting in a final selection of 42 participants. The distribution of participants is as follows: 6 CTOs or equivalent, 3 CIOs or equivalent, 13 IT Directors or equivalent, 16 IT Project Managers, 2 IT Security Managers, and 2 IT Auditors, as shown in Figure 2.

Participant demographics included 31 Southeast Asians, 3 Indians, 4 Chinese, 1 Japanese, 2 Thai working in the U.S., and 1 Russian, highlighting a diverse range of perspectives. All were experienced in IT systems related to unstructured data. The data collection process involved:

- 1) Initial Contact via various channels.
- 2) Interviews using a 16-question framework from "The Framework of Extracting Unstructured Usage for Big Data Platform,"
- 3) Data Processing with a 6-point scoring system,
- 4) Human Review of qualitative responses, and
- 5) Data Verification with participants. This process is depicted in Figure 3.

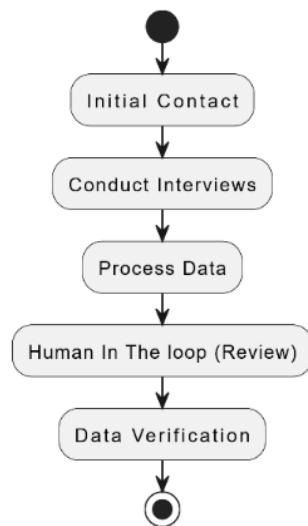


Figure 3: Data collection process

B. Data Separation with RAG Technique.

The collected data from 42 participants, comprising responses to 16 questions, covered various aspects including physical characteristics, Time, Priority, Resources, and Stakeholder dimensions. To process this data, interviews were transcribed and refined using GPT, while written responses were converted to text. The resulting text was then segmented and mapped to relevant concepts from "The Framework of Extracting Unstructured Usage for Big Data Platform" using

Retrieval Augmented Generation (RAG) with GPT-4. This approach minimized hallucinations and enhanced accuracy in the data processing [15]. The context-tuned planner, based on the work of Anantha et al. (2024), achieved a notable AST-based Plan Accuracy of 85.24%, while significantly reducing hallucinations to 0.93%. This improvement underscores the importance of context integration in enhancing the accuracy and reliability of the planning process for unstructured data analysis. The Context-tuned Upper Bound was selected for performing RAG, ensuring that sentence segmentation remained within the contextual framework. This choice was crucial for preventing excessive hallucinations and ensuring more accurate and contextually relevant segmentation, which is essential for analyzing unstructured data in IT projects [15]. As Table 4 illustrates:

Table 4: End-to-end Planner Evaluation

Setting	AST-based Plan Acc	Exact Match	Hallucination
Lower Bound	43.77	39.45	2.59
RAG-based Planner	76.39	58.12	1.76
Context-tuned RAG Planner	85.24	67.33	0.93
Upper Bound	91.47	72.65	0.85
Context-tuned Upper Bound	91.62	72.84	0.53

C. Application of AI Techniques Zero-shot Text Classification

The Zero-Shot Text Classification approach uses Python tools to interpret context, with interviews conducted mainly in Thai and English. The code was adapted for both languages, using facebook/bart-large-mnli for English and joeddav/xlm-roberta-large-xnli for Thai. The joeddav/xlm-roberta-large-xnli model, a multilingual extension of RoBERTa, supports Thai and is trained for cross-lingual inference (XNLI). In contrast, facebook/bart-large-mnli is a monolingual model

designed for English, utilizing the BART architecture for text generation and transformation.

Trained for MNLI (Multi-Genre Natural Language Inference), which involves understanding linguistic inference in various English genres. Reasons for Choosing joeddav/xlm-roberta-large-xnli for Thai:

The joeddav/xlm-roberta-large-xnli model, trained for Thai, outperforms the facebook/bart-large-mnli model, which only supports English, by directly processing Thai text and reducing translation errors. For example, in response to "What documents need to be scanned and how is the workflow prioritized, including resources and time management?" the model effectively applies Zero-Shot Text Classification, crossing Object x Property and incorporating Project Management dimensions like Time, Priority, Resources, and Stakeholder." The Thai response is: "ใบแจ้งหนี้ที่เราต้องการแปลงเป็นรูปแบบดิจิทัล ซึ่งเป็นสิ่งสำคัญสำหรับระบบการจัดการเอกสารของเรา จะได้รับทุกวันก่อน 12:00 ผ่านไปรษณีย์ไทย ซึ่งจะมาปนกับจดหมายอื่นๆ เราจำเป็นต้องคัดแยกเพื่อนำไปสแกนเอกสารและ OCR เพื่อให้สามารถประมวลผลและจัดการข้อมูลได้อย่างมีประสิทธิภาพภายใน 2 ชั่วโมง"

This translates to English: "The invoices we need to digitize, which are crucial for our document management system, are received daily before 12:00 PM via Thai Post, mixed with other mail. We need to sort them for scanning and optical character recognition to efficiently process and manage the data within 2 hours."

When the example context is run using the Zero-Shot model joeddav/xlm-roberta-large-xnli, the results are as shown in Figure 4. For comparison, when the same context translated into English is run using the facebook/bart-large-mnli model, the results differ as shown in Table 5.

```
from transformers import pipeline

# Load the Zero-Shot Classification Pipeline with a multilingual model
classifier = pipeline("zero-shot-classification", model="joeddav/xlm-roberta-large-xnli")

# The sentence to classify
sentence = "ใบแจ้งหนี้ที่เราต้องการแปลงเป็นรูปแบบดิจิทัล ซึ่งเป็นสิ่งสำคัญสำหรับระบบการจัดการเอกสารของเรา จะได้รับทุกวันก่อน 12:00 ผ่านไปรษณีย์ไทย ซึ่งจะมาปนกับจดหมายอื่นๆ เราจำเป็นต้องคัดแยกเพื่อนำไปสแกนเอกสารและ OCR เพื่อให้สามารถประมวลผลและจัดการข้อมูลได้อย่างมีประสิทธิภาพภายใน 2 ชั่วโมง"

# Labels to check for relevance
labels = ["Object", "Property", "Time", "Priority", "Resources", "Stakeholder"]

# Perform classification for each label and check results
threshold = 0.10
for label in labels:
    result = classifier(sentence, [label])
    if result['scores'][0] >= threshold:
        print(f"Relevant to {label}")
    else:
        print(f"Not relevant to {label}")
```

Some weights of the model checkpoint at joeddav/xlm-roberta-large-xnli were not used when
- This IS expected if you are initializing XLMRobertaForSequenceClassification from the ch
- This IS NOT expected if you are initializing XLMRobertaForSequenceClassification from th
Not relevant to Object
Relevant to Property
Relevant to Time
Relevant to Priority
Relevant to Resources
Not relevant to Stakeholder

Figure 4: Zero-Shot model joeddav/xlm-roberta-large-xnli

Table 5: Compare for Thai response and English

	joeddav/xlm-roberta-large-xnli	facebook/bart-large-mnli
Object	Not relevant	Relevant
Property	Relevant	Relevant
Time	Relevant	Relevant
Priority	Relevant	Relevant
Resources	Relevant	Relevant
Stakeholder	Not relevant	Relevant

This table compares the performance of joeddav/xlm-roberta-large-xnli and facebook/bart-large-mnli models in classifying Thai and English text respectively. The results show that the Thai model (joeddav/xlm-roberta-large-xnli) performs better in identifying 'Property', 'Time', and 'Resources' categories, while the English model (facebook/bart-large-mnli) excels in 'Object' and 'Stakeholder' categories. This difference in performance highlights the importance of language-specific models in multilingual text classification tasks.

D. Generative AI Text Classification

To ensure accurate classification by ChatGPT, the PARTS Framework (Persona, Action, Result, Target, and Style) was employed for prompting, providing a structured approach to generating precise and context-appropriate prompts for AI models. This framework defines the AI's

role, specifies the task, outlines expected outcomes, identifies the intended audience, and determines the response tone, leading to improved classification results. A portion of the prompt used is shown in Figure 5, with the results summarized in Table 6.

```

**Purpose:**
Read and understand the text, then determine if it
relates to the given set of keywords.

**Audience:**
Computer program processing relevance as Y or N for
further analysis.

**Requirements:**
Accepts text in both Thai and English, written formally.

**Tone:**
Polite

**Scope:**
1. First, ask for the "text" that needs to be analyzed
for relationships.
2. Then, ask for the set of keywords for classification,
explaining that they will be used to check relevance.
3. Once both the text and the set of keywords are
obtained, understand their context.
4. If they are related, explain how they are related.

```

Figure 5: Sample Prompt PARTS Framework

Table 6: Results of Generative AI Text Classification

Keyword	Relevance
Object	Relevant
Property	Not relevant
Time	Relevant
Priority	Relevant
Resources	Relevant
Stakeholder	Not relevant

The results demonstrate the Generative AI model's effectiveness in identifying certain categories ('Object', 'Time', 'Priority', and 'Resources') while showing limitations in others ('Property' and 'Stakeholder'). This pattern suggests that the Generative AI approach may be more suitable for tasks focusing on temporal and resource-related aspects of data in IT projects.

E. Comparative Analysis and Validation and Human-in-the-Loop.

When both techniques, Zero-Shot and Generative AI, were used and validated by the respondents, the

results were obtained as shown in Table 7. Reduce the size of the table, "Relevant" is denoted as "Y" and "Not relevant" as "N".

Table 7: Comparative Results

	joeddav	facebook	Generative AI	Human In The Loop
Object	Y	Y	Y	Y
Property	Y	Y	N	Y
Time	Y	Y	Y	Y
Priority	N	Y	Y	N
Resources	Y	Y	Y	Y
Stakeholder	Y	Y	N	Y

This comparative analysis provides insights into the strengths and limitations of each AI approach relative to human judgment. The results suggest that while AI models show promising performance, there are still areas where they differ from human classification, particularly in categories like 'Priority' and 'Stakeholder'.

F. Summary of Key Findings

This analysis used Human In The Loop data to benchmark the reliability of three models: joeddav/xlm-roberta-large-xnli, facebook/bart-large-mnli, and a Generative model. A Proportion Test assessed each model's reliability against a standard. The proportion of matches was calculated, followed by a weighted average. Z-scores were determined and compared to a Z-critical value of 1.96 for a 95% confidence level. With a sample size of 42 participants answering 16 questions (672 responses), results are summarized in Table 7.

G. Summary of Comparative Analysis and Validation

1) joeddav/xlm-roberta-large-xnli shows high reliability, as its Z-value is 3.41, exceeding the Z-critical value (1.96).

2) Generative also shows high reliability, as its Z-value is -2.68, which is significant at a 95% confidence level (in the negative direction indicating significant deviation).

3) facebook/bart-large-mnli is not statistically significant, with a Z-value of -0.73, which is below the Z-critical value (1.96).

Table 8: Z-Test Hypothesis Results

Model	Matches	Mismatches	Matches (%)	Mismatches (%)	Z-value	Significant
joeddav	624	48	93%	7%	3.41	Yes
facebook	590	82	88%	12%	-0.73	No
Generativ	574	98	85%	15%	-2.68	Yes

The analysis shows that the joeddav/xlm-roberta-large-xnli and Generative models are more reliable than the Human In The Loop standard, while the facebook/bart-large-mnli model falls short in comparison.

These results provide quantitative evidence for the relative strengths of each approach. The joeddav model shows the highest reliability, suggesting its potential for multilingual applications in IT projects. The Generative model's significant result, albeit in the negative direction, indicates its unique approach to classification tasks.

H. Integrated Analysis of Results

The collective analysis of Tables 4-8 provides a comprehensive view of the performance and reliability of different AI techniques in transforming unstructured qualitative data into quantitative insights within IT projects. The RAG technique (Table 4) demonstrates the importance of context in improving accuracy and reducing errors. The comparison of language-specific models (Table 5) highlights the nuanced differences in processing Thai and English text, which is crucial for multilingual IT environments. The Generative AI results (Table 6) show its strength in certain categories, particularly those related to time and resources, which are often critical in IT project management. The human-in-the-loop validation (Table 7) offers insights into how well these AI models align with human judgment, an important factor in practical applications. Finally, the statistical analysis (Table 8) provides a quantitative basis for comparing the reliability of these different approaches.

Overall, these results suggest that while each method has its strengths, the joeddav/xlm-roberta-large-xnli model demonstrates the highest overall reliability for this specific task in IT projects. However, the strong performance of the Generative AI model in certain categories indicates its potential for specialized applications within IT project management, particularly in areas focusing on temporal and resource-related data. These findings have significant implications for choosing appropriate AI techniques for different types of unstructured data in various IT project contexts.

V. DISCUSSION

The findings from this study highlight the efficacy and reliability of different AI models in transforming qualitative data into quantitative insights. The Zero-Shot Text Classification and Generative AI Text Classification techniques were rigorously evaluated



using a sample of 42 participants who provided 672 responses. These responses were analyzed and validated against Human In The Loop data to benchmark the performance of three AI models: joeddav/xlm-roberta-large-xnli, facebook/bart-large-mnli, and Generative.

A. Zero-Shot Text Classification

The Zero-Shot Text Classification approach, especially with the joeddav/xlm-roberta-large-xnli model, showed high reliability. With a Z-value of 3.41, well above the critical 1.96, this model effectively handled Thai text, which was the majority of the data. Its multilingual capabilities make it a robust tool for organizations dealing with multilingual unstructured data, offering reliable classifications across different languages and contexts.

B. Generative AI Text Classification.

The Generative AI Text Classification technique demonstrated strong performance, with a Z-value of -2.68. Using the PARTS Framework for prompting, it generated accurate classifications validated by human respondents. This model's ability to contextualize data with minimal pre-training indicates its potential for handling complex datasets. Its performance suggests that generative approaches offer flexible, adaptive solutions for real-time data classification, suitable for dynamic environments.

C. Comparison & Implications

Comparing the AI techniques reveals key strengths: joeddav/xlm-roberta-large-xnli excels in structured, multilingual tasks with a high reliability, while the Generative AI model is more adaptable for dynamic environments. The facebook/bart-large-mnli model, with a Z-value of -0.73, lacks the robustness of the other models, suggesting limited utility and potential areas for future enhancement.

D. Practical Applications.

These findings have practical implications for organizations handling large volumes of unstructured data. The joeddav/xlm-roberta-large-xnli model can be used in multilingual document management, while Generative AI can dynamically classify customer queries in service platforms. Integrating these models can improve data processing, decision-making, and productivity.

VI. CONCLUSION

This study highlights the vital role of advanced AI in transforming qualitative data into quantitative insights. The evaluation of joeddav/xlm-roberta-large-xnli, facebook/bart-large-mnli, and Generative AI models demonstrates that Zero-Shot and Generative AI approaches offer reliable, scalable solutions. With high Z-values, joeddav/xlm-roberta-large-xnli excels in multilingual support, while Generative AI shines in adaptability and contextual understanding.

These models enhance data classification accuracy and reduce the need for extensive pre-training, making them invaluable for efficient unstructured data management. Their integration into data systems promises significant improvements in processing accuracy and efficiency. Future research should refine these models and explore new applications, fully harnessing AI's potential in managing unstructured data.

REFERENCES

- [1] H.-J. Kong, "Managing unstructured big data in healthcare system," *Healthc. Inform. Res.*, vol. 25, no. 1, pp. 1–2, 2019.
- [2] H. Abburi, M. Suesserman, N. Pudota, B. Veeramani, E. Bowen, and S. Bhattacharya, "Generative AI text classification using ensemble LLM approaches," in *Proc. Iberian Lang. Eval. Forum*, Jaén, Spain, Sep. 2023. [Online]. Available: <https://ceur-ws.org/Vol-3496/autextification-paper14.pdf>

- [3] C. Han, H. Pei, X. Du, and H. Ji, "Zero-shot classification by logical reasoning on natural language explanations," in *Proc. 61st Annu. Meeting Assoc. Comput. Linguistics*, Toronto, Canada, Jul. 2023. pp. 8967–8981.
- [4] J. Zhang, P. Lertvittayakumjorn, and Y. Guo, "Integrating semantic knowledge to tackle zero-shot text classification," in *Proc. Conf. North Amer. Chapter Assoc. Comput. Linguistics: Human Lang. Technologies*, Minneapolis, MN, USA, Jun. 2019, pp. 1031–1040.
- [5] Forbes Thailand. "Unstructured Data: A treasure trove waiting to be unlocked." (in Thai), FORBESTHAILAND.com. <https://forbesthailand.com/commentaries/Insights/unstructured-data-ข้อมูลที่รอการปลด> (accessed Aug. 1, 2024).
- [6] Q. Ye et al., "Zero-shot text classification via reinforced self-training," in *Proc. 58th Annu. Meeting Assoc. Comput. Linguistics*, Seattle, WA, USA, Jul. 2020, pp. 3014–3024.
- [7] W. Yin, J. Hay, and D. Roth, "Benchmarking zero-shot text classification: Datasets, evaluation and entailment approach," in *Proc. Conf. Empirical Methods in Natural Lang. Process. and 9th Int. Joint Conf. Natural Lang. Process. (EMNLP-IJCNLP)*, Hong Kong, China, Nov. 2019, pp. 3914–3923.
- [8] T. Brown et al., "Language models are few-shot learners," in *Proc. 34th Conf. Neural Inf. Process. Syst. (NeurIPS)*, Vancouver, Canada, Dec. 2020, pp.1–25.
- [9] K. J. Srnka and S. T. Koeszegi, "From words to numbers: How to transform qualitative data into meaningful quantitative results," *Schmalenbach Bus. Rev.*, vol. 59, pp. 29–57, 2007.
- [10] H. Abi Akl, "A ML-LLM pairing for better code comment classification," in *Proc. 15th Meeting Forum Inf. Retrieval Eval. (FIRE)*, Panjim, India, Dec. 2023.
- [11] T.-l. Chasupa and W. Paireekreng, "The framework of extracting unstructured usage for big data platform," in *Proc. 2nd Int. Conf. Big Data Analytics and Pract. (IBDAP)*, Bangkok, Thailand, Aug. 2021, pp. 90–94.
- [12] Talance. "Let's take a look! What are the responsibilities of IT positions?." (in Thai), TALANCE.tech <https://www.talance.tech/blog/it-job-responsibility/> (accessed Aug. 1, 2024).
- [13] G. Jaimovitch-López, C. Ferri, J. Hernández Orallo, F. Martínez-Plumed, and M. J. Ramírez-Quintana, "Can language models automate data wrangling?," *Mach. Learn.*, vol. 112, pp. 2053–2082, 2023.
- [14] H. Ahmad and H. Halim, "Determining sample size for research activities: The case of organizational research," *Selangor Bus. Rev.*, vol. 2, no. 1, pp. 20–34, 2017.
- [15] R. Anantha, T. Bethi, D. Vodianik, and S. Chappidi, "Context tuning for retrieval augmented generation," in *Proc. 1st Workshop on Uncertainty-Aware NLP (UncertainNLP)*, St. Julian's, Malta, Mar. 2024, pp. 15–22.

คำแนะนำสำหรับผู้เขียนบทความเพื่อลงตีพิมพ์

วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล เป็นวารสารวิชาการทางด้านวิศวกรรมศาสตร์และเทคโนโลยี ของสถาบันเทคโนโลยีไทย-ญี่ปุ่น บทความที่น่าสนใจจะต้องพิมพ์เป็นภาษาไทย หรือภาษาอังกฤษตามรูปแบบที่กำหนด และพร้อมที่จะนำไปตีพิมพ์ได้ทันที การเสนอบทความเพื่อพิจารณาตีพิมพ์ในวารสาร มีรายละเอียดดังนี้

1. หลักเกณฑ์การพิจารณาบทความเพื่อตีพิมพ์

1.1 เป็นบทความที่ไม่ได้อยู่ระหว่างการพิจารณาตีพิมพ์ หรือไม่ได้อยู่ระหว่างการพิจารณาของสื่อสิ่งพิมพ์อื่น ๆ และไม่เคยตีพิมพ์ในวารสารหรือรายงานการสืบเนื่องจากการประชุมวิชาการใดมาก่อนทั้งในประเทศและต่างประเทศ หากตรวจสอบพบว่ามีการตีพิมพ์ซ้ำซ้อน ถือเป็นความรับผิดชอบของผู้เขียนแต่เพียงผู้เดียว

1.2 เป็นบทความที่แสดงให้เห็นถึงความคิดริเริ่มสร้างสรรค์ มีคุณค่าทางวิชาการ มีความสมบูรณ์ของเนื้อหา และมีความถูกต้องตามหลักวิชาการ

1.3 บทความที่ได้รับการตีพิมพ์จะต้องผ่านการประเมินคุณภาพจากผู้ทรงคุณวุฒิ (Peer reviewer) อย่างน้อย 2 ท่าน ต่อบทความ ซึ่งผู้ทรงคุณวุฒิอาจให้ผู้เขียนแก้ไขเพิ่มเติมหรือปรับปรุงบทความให้เหมาะสมยิ่งขึ้น

1.4 กองบรรณาธิการขอสงวนสิทธิ์ในการตรวจแก้ไขรูปแบบบทความที่ส่งมาตีพิมพ์ตามที่เห็นสมควร

1.5 บทความ ข้อความ ภาพประกอบ และตารางประกอบ ที่ตีพิมพ์ลงวารสารเป็นความคิดเห็นส่วนตัวของผู้เขียน กองบรรณาธิการไม่มีส่วนรับผิดชอบใด ๆ ถือเป็นความรับผิดชอบของผู้เขียนแต่เพียงผู้เดียว

1.6 ต้องเป็นบทความที่ไม่ละเมิดลิขสิทธิ์ ไม่ลอกเลียน หรือคัดลอกข้อความของผู้อื่นโดยไม่ได้รับอนุญาต

1.7 หากเป็นงานแปลหรือเรียบเรียงจากภาษาต่างประเทศ ต้องมีหลักฐานการอนุญาตให้ตีพิมพ์เป็นลายลักษณ์อักษรจากเจ้าของลิขสิทธิ์

1.8 ต้องมีการอ้างอิงที่ถูกต้อง เหมาะสมก่อนการตีพิมพ์ ซึ่งเป็นความรับผิดชอบของเจ้าของผลงาน

1.9 บทความที่ส่งถึงกองบรรณาธิการ ขอสงวนสิทธิ์ที่จะไม่ส่งคืนผู้เขียน

2. รูปแบบการกลั่นกรองบทความก่อนลงตีพิมพ์ (Peer-review)

ในการประเมินบทความโดยผู้ทรงคุณวุฒิเป็นการประเมินแบบ Double-blind peer review คือ ผู้ทรงคุณวุฒิไม่ทราบชื่อและรายละเอียดของผู้เขียนบทความ และผู้เขียนบทความไม่ทราบชื่อและรายละเอียดของผู้ทรงคุณวุฒิ

3. ประเภทของบทความที่รับพิจารณาลงตีพิมพ์

นิพนธ์ต้นฉบับต้องเป็นบทความวิจัย ประกอบด้วย บทคัดย่อ บทนำ วัตถุประสงค์ของการวิจัย วิธีดำเนินการวิจัย ผลการวิจัยและอภิปรายผล สรุปผล และเอกสารอ้างอิง*

หมายเหตุ : บทความภาษาไทยต้องมีบทคัดย่อภาษาไทยและภาษาอังกฤษ โดยให้หน้าบทคัดย่อภาษาไทยอยู่ก่อนหน้าบทคัดย่อภาษาอังกฤษ สำหรับบทความภาษาอังกฤษไม่ต้องมีบทคัดย่อภาษาไทย

* เอกสารอ้างอิง เป็นการบอกรายการแหล่งอ้างอิงที่มีการอ้างอิงในเนื้อหาของงานเขียน

4. จริยธรรมในการตีพิมพ์ผลงานวิจัยในวารสารวิชาการ

วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล ได้คำนึงถึงจริยธรรมในการตีพิมพ์บทความ โดยจริยธรรมและบทบาทหน้าที่ของผู้เกี่ยวข้อง มีดังนี้

จริยธรรมและหน้าที่ของผู้เขียน

1. ผู้เขียนต้องเขียนบทความให้เป็นไปตามรูปแบบที่วารสารกำหนดไว้ในคำแนะนำสำหรับผู้เขียน
2. หากมีการนำข้อมูลของผู้อื่นหรือข้อมูลของผู้เขียนที่เคยตีพิมพ์ในวารสารฉบับอื่นมาใช้ ผู้เขียนต้องอ้างอิงแหล่งที่มาของข้อมูลนั้น โดยไม่ทำให้ผู้อื่นเข้าใจผิดว่าข้อมูลนั้นเป็นผลงานใหม่ของผู้เขียน
3. ผู้เขียนต้องไม่ดัดแปลงหรือบิดเบือนข้อมูล
4. ผู้เขียนต้องระบุแหล่งทุนที่สนับสนุนในการทำวิจัย (ถ้ามี)
5. ผู้เขียนต้องเปิดเผยข้อมูลเกี่ยวกับผลประโยชน์ทับซ้อนอย่างชัดเจน (ถ้ามี)

จริยธรรมและหน้าที่ของบรรณาธิการ

1. บรรณาธิการต้องดำเนินการเผยแพร่วารสารให้ตรงตามเวลา
2. บรรณาธิการต้องคัดเลือกบทความ โดยพิจารณาจากคุณภาพและความสอดคล้องของเนื้อหาบทความกับขอบเขตของวารสาร
3. บรรณาธิการต้องดำเนินการประเมินบทความอย่างเป็นธรรม ไม่ปฏิเสธการตีพิมพ์บทความโดยใช้อคติ
4. บรรณาธิการต้องไม่เปิดเผยข้อมูลของผู้เขียนบทความ และผู้ประเมินบทความแก่บุคคลอื่นที่ไม่เกี่ยวข้อง
5. บรรณาธิการต้องไม่มีผลประโยชน์ทับซ้อนกับผู้เขียน และผู้ประเมินบทความ
6. บรรณาธิการต้องใช้โปรแกรมในการตรวจสอบการคัดลอกผลงาน เพื่อป้องกันการตีพิมพ์ผลงานซึ่งคัดลอกมาจากผลงานผู้อื่น หากพบการคัดลอกผลงานของผู้อื่น บรรณาธิการต้องหยุดกระบวนการพิจารณาบทความทันที และติดต่อผู้เขียนเพื่อขอคำชี้แจง

จริยธรรมและหน้าที่ของผู้ประเมินบทความ

1. ผู้ประเมินบทความ ควรพิจารณาตอบรับการประเมินเฉพาะบทความที่สอดคล้องกับความเชี่ยวชาญของตนเองเท่านั้น เพื่อให้บทความที่ตีพิมพ์มีคุณภาพ
2. ผู้ประเมินบทความควรประเมินบทความให้แล้วเสร็จตามระยะเวลาที่กำหนด เพื่อไม่ให้บทความที่ผ่านการพิจารณาตีพิมพ์ล่าช้า
3. ผู้ประเมินควรประเมินบทความโดยให้ข้อเสนอแนะตามหลักวิชาการเท่านั้น ไม่ควรใช้ความคิดเห็นส่วนตัวที่ไม่มีเหตุผลหรือไม่มีข้อมูลรองรับ
4. ผู้ประเมินควรปฏิเสธการประเมินบทความ หากเห็นว่าตนเองอาจมีผลประโยชน์ทับซ้อนกับผู้เขียน
5. ผู้ประเมินต้องไม่เปิดเผยเนื้อหาภายในบทความให้ผู้ที่ไม่เกี่ยวข้องทราบ
6. หากผู้ประเมินเห็นว่ามีความเห็นอื่นที่ผู้เขียนไม่ได้กล่าวอ้างถึง แต่เป็นบทความที่สำคัญและเกี่ยวข้องกับบทความ ผู้ประเมินควรแจ้งให้ผู้เขียนกล่าวอ้างถึงบทความนั้น

5. ข้อกำหนดการจัดพิมพ์ต้นฉบับบทความ

ผู้เขียนต้องจัดพิมพ์บทความตามข้อกำหนดเพื่อให้มีรูปแบบการตีพิมพ์เป็นมาตรฐานแบบเดียวกัน ดังนี้

5.1 ขนาดของกระดาษ ให้ใช้ขนาด A4

5.2 กรอบของข้อความ ระยะห่างของขอบกระดาษ

ด้านบน 2.5 ซม. (0.98")

ด้านล่าง 2 ซม. (0.79")

ด้านซ้าย 2 ซม. (0.79")

ด้านขวา 2 ซม. (0.79")

5.3 ระยะห่างระหว่างบรรทัด หนึ่งช่วงบรรทัดของเครื่องคอมพิวเตอร์ (Single)

5.4 ตัวอักษร รูปแบบของตัวอักษรให้ใช้ TH Sarabun New

5.5 รายละเอียดต่าง ๆ ของบทความ กำหนดดังนี้

ชื่อเรื่อง (Title) ขนาด 22 ตัวหนา กำหนดกึ่งกลาง

ชื่อผู้เขียน (Author) ขนาด 16 ตัวธรรมดา กำหนดกึ่งกลาง ไม่ต้องใส่คำนำหน้า

สังกัด (Affiliation) ขนาด 14 ตัวเอน กำหนดกึ่งกลาง

E-mail ผู้ประพันธ์บรรณกิจ (Corresponding Author) ขนาด 12 ตัวธรรมดา กำหนดกึ่งกลาง

บทคัดย่อ (Abstract) จัดรูปแบบการพิมพ์เป็นแบบ 1 คอลัมน์ ชื่อหัวข้อ ขนาด 14 **ตัวหนาและเอน** กำหนดกึ่งกลาง ข้อความในบทคัดย่อ ขนาด 14 ตัวธรรมดา

คำสำคัญ (Keywords) ให้ใส่คำสำคัญ 3–5 คำ ซึ่งเกี่ยวข้องกับบทความที่นำเสนอ โดยให้จัดพิมพ์ใต้บทคัดย่อ เว้น 1 บรรทัดด้วยขนาด 12 คำสำคัญขนาด 14 **ตัวหนาและเอน** กำหนดชิดซ้าย ข้อความในคำสำคัญ ขนาด 14 ตัวธรรมดา

รูปแบบการพิมพ์เนื้อหาของบทความ

- รูปแบบการพิมพ์เป็นแบบ 2 คอลัมน์ แต่ละคอลัมน์ กว้าง 8.2 ซม. (3.23") ระยะห่างระหว่างคอลัมน์ 0.61 ซม. (0.24")

- หัวข้อหลัก ประกอบด้วย บทนำ (INTRODUCTION) วัตถุประสงค์ของการวิจัย (OBJECTIVES) วิธีดำเนินการวิจัย (METHODS) ผลการวิจัยและอภิปรายผล (RESULTS AND DISCUSSION) สรุปผล (CONCLUSIONS) เอกสารอ้างอิง (REFERENCES) ขนาด 14 ตัวธรรมดา กำหนดกึ่งกลาง และมีเลขกำกับ เช่น 1) บทนำ หรือ I. INTRODUCTION เป็นต้น

- หัวข้อรอง ระดับที่ 1 ขนาด 14 ตัวเอน กำหนดชิดซ้าย เมื่อขึ้นหัวข้อระดับเดียวกันหรือมากกว่าในลำดับถัดไป ให้เคาะเว้น 1 บรรทัดหลังจบเนื้อหา

- หัวข้อรอง ระดับที่ 2 ขนาด 14 ตัวเอน กำหนดชิดซ้ายและเลื่อนเข้ามา 0.5 ซม.

- เนื้อเรื่อง ขนาด 14 ตัวธรรมดา

- ชื่อตาราง ขนาด 12 ตัวธรรมดา กำหนดกึ่งกลาง และใส่ชื่อเหนือตาราง

- หัวข้อในตาราง ขนาด 12 ตัวหนา กำหนดกึ่งกลาง เนื้อหาในตาราง ขนาด 12 ตัวธรรมดา

- ชื่อภาพประกอบ ขนาด 12 ตัวธรรมดา กำหนดกึ่งกลาง และใส่ชื่อใต้ภาพ

- เนื้อหาในภาพประกอบ ขนาด 12 ตัวธรรมดา

วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล

ปีที่ 13 ฉบับที่ 2 กรกฎาคม – ธันวาคม 2568

5.6 เอกสารอ้างอิง

1. การอ้างอิงในเนื้อหาบทความใช้การอ้างอิงแบบตัวเลข ตามมาตรฐานสากล โดยใช้หมายเลขในเครื่องหมายก้ามปู [] และเรียงลำดับการอ้างอิงตามเนื้อหา โดยมีตัวอย่างการเขียน เช่น [1] หรือ [2] หรือ [1], [2] หรือ [1], [3]–[8] หรือ [9], [10], [15], [16] หากมีการอ้างอิงซ้ำบทความเดิมให้ใช้หมายเลขเดิม ตัวอย่างเช่น by Brown [4], [5]; as mentioned earlier [2], [4]–[7], [9]; Smith [4] and Brown and Jones [5]; Wood *et al.* [7]
2. รูปแบบของชื่อหัวข้อใช้รูปแบบตัวอักษร TH Sarabun New ขนาด 14 ตัวธรรมดา ในเนื้อหาขนาด 12 ตัวธรรมดา
3. การอ้างอิงท้ายบทความ จะต้องเรียงตามลำดับบทความที่เขียนอ้างอิงในเนื้อเรื่อง และใช้การอ้างอิงตามรูปแบบการอ้างอิง IEEE V 11.12.2018 ซึ่งผู้เขียนสามารถศึกษาวิธีการเขียนเอกสารอ้างอิงตามรูปแบบที่กำหนดได้ที่เว็บไซต์ <https://ieeauthorcenter.ieee.org/wp-content/uploads/IEEE-Reference-Guide.pdf> โดยจะต้องเขียนเป็นภาษาอังกฤษเท่านั้น หากบทความอ้างอิงมาจากบทความภาษาไทย ต้องแปลเป็นภาษาอังกฤษให้ถูกต้อง
4. กรณีที่เอกสารที่นำมาอ้างอิงเขียนเป็นภาษาไทยให้เติมคำว่า “(in Thai)” เข้าไปในเอกสารอ้างอิง ดังเช่นตัวอย่างต่อไปนี้

ตัวอย่างรูปแบบการเขียนและการแปลเอกสารอ้างอิงภาษาไทยเป็นภาษาอังกฤษ

ตัวอย่างที่ 1 การอ้างอิงจากหนังสือ (Books)

Basic Format:

- [Number] J. K. Author, *Title of His Published Book*, xth ed. City of Publisher, State (only U.S.), Country: Abbrev. of Publisher, year.
- [Number] J. K. Author, “Title of chapter in the book,” in *Title of His Published Book*, xth ed. City of Publisher, State (only U.S.), Country: Abbrev. of Publisher, year, ch. x, sec. x, pp. xxx–xxx.

Examples:

- [1] V. Rijiravanich, *Work Study: Principles and Case Studies*, 4th ed. Bangkok, Thailand: Chulalongkorn University Press (in Thai), 2005.
วันชัย ริจิรวณิช, *การศึกษาการทำงาน: หลักการและกรณีศึกษา*, พิมพ์ครั้งที่ 4, กรุงเทพฯ, ประเทศไทย: สำนักพิมพ์จุฬาลงกรณ์มหาวิทยาลัย, 2548.
- [2] L. Edelstein-Keshet, *Mathematical Models in Biology*. New York, NY, USA: Random House, 1998.
- [3] K. J. Roodbergen, “Storage assignment for order picking in multiple-block warehouses,” in *Warehousing in the Global Supply Chain: Advanced Models, Tools and Applications for Storage Systems*. London, U.K.: Springer, 2012, pp. 139–155.

ตัวอย่างที่ 2 การอ้างอิงจากวารสาร (Periodicals)

Basic Format:

[Number] J. K. Author, "Name of paper," *Abbrev. Title of Periodical*, vol. x, no. x, pp. xxx-xxx, Abbrev. Month. year.

[Number] J. K. Author, "Name of paper," *Abbrev. Title of Periodical*, vol. x, no. x, pp. xxx-xxx, Abbrev. Month. year, doi: xxx.

Examples:

[4] N. Dechampai and K. Sethanan, "An application of lean manufacturing system in the textile of lean manufacturing system in the textile and garment industry case study: Wacoal Kabinburi Co., Ltd," (in Thai), *MBA-KKU J.*, vol. 7, no. 2, pp. 13-27, 2014.

นิวัฒน์ เดชอำไพ และกาญจนา เศรษฐนันท์, "การเพิ่มประสิทธิภาพกระบวนการผลิตชุดชั้นในสตรีโดยประยุกต์ใช้แนวคิดการผลิตแบบลีน," *วารสารวิทยาลัยบัณฑิตศึกษากิจการมหาวิทยาลัยขอนแก่น*, ปีที่ 7, ฉบับที่ 2, หน้า 13-27, 2557.

[5] S.-F. Wang, H.-P. Chen, Y. Ku, and M.-X. Zhong, "Voltage-mode multifunction biquad filter and its application as fully-uncoupled quadrature oscillator based on current-feedback operational amplifiers," *Sensors*, vol. 20, no. 22, p. 6681, Nov. 2020, doi: 10.3390/s20226681.

[6] M. Dwiyaniti *et al.*, "Extremely high surface area of activated carbon originated from sugarcane bagasse," in *IOP Conf. Ser.: Mater. Sci. Eng.*, vol. 909, no. 1, 2020, pp. 1-8.

[7] M. Yoshiki, Y. Fujita, A. Kawamura, and H. Arai, "Instability of plates with holes (1st report)," (in Japanese), *J. Zosen Kiokai*, vol. 1967, no. 122, pp. 137-145, 1967.

[8] F. de Oliveira Santini, "Clockwise versus counterclockwise turning bias: Moderation effects of foot traffic and cognitive experience on visual attention," *J. Retail. Consum. Serv.*, vol. 67, 2022, Art. no. 102965.

ตัวอย่างที่ 3 การอ้างอิงจากวิทยานิพนธ์ (Theses and Dissertations)

Basic Format:

[Number] J. K. Author, "Title of thesis," M.S. thesis / Ph.D. dissertation, Abbrev. Dept., Abbrev. Univ., City of Univ., Abbrev. State (only U.S.), year.

Examples:

[9] N. Kawasaki, "Parametric study of thermal and chemical nonequilibrium nozzle flow," M.S. thesis, Dept. Electron. Eng., Osaka Univ., Osaka, Japan, 1993.

[10] Y. Sornsa and P. Wanrerak, "Buckling of rectangular plates with a central square hole," (in Thai), B.S. thesis, Dept. Mech. Eng., Burapha Univ., Chonburi, Thailand, 2007.

วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล

ปีที่ 13 ฉบับที่ 2 กรกฎาคม – ธันวาคม 2568

ตัวอย่างที่ 4 การอ้างอิงจากการประชุมทางวิชาการ (Conferences and Conference Proceedings)

Basic Format:

[Number] J. K. Author, "Title of paper," in *Abbrev. Name of Conf.*, City, State (only U.S.), Country, Abbrev. Month. year, pp. xxx–xxx.

Examples:

[11] N. Kriengkarakot, P. Kriengkarakot, S. Duan, P. Thung, and W. Piromsuk, "Repair work reduction in sewing process of the apparel factory," (in Thai), in *Proc. 10th Ubon Ratchathani Univ. Nat. Res. Conf.*, Ubon Ratchathani, Thailand, Jul. 2016, pp. 87–96.

นุชสรา เกรียงกรกฎ, ปรีชา เกรียงกรกฎ, สกาวเดือน พรหมทุ่ง และวิจิตรา ภิรมย์สุข, "งานวิจัยการลดชิ้นส่วนงานซ่อมในขั้นตอนการเย็บของโรงงานผลิตเสื้อผ้าสำเร็จรูป," *การประชุมวิชาการระดับชาติ มอช. วิจัยครั้งที่ 10*, อุบลราชธานี, ประเทศไทย, 7-8 กรกฎาคม, 2559, หน้า 87–96.

[12] J. Lingad, S. Karimi, and J. Yin, "Location extraction from disaster-related microblogs," in *Proc. 22nd Int. Conf. World Wide Web*, Rio de Janeiro, Brazil, May 2013, pp. 1017–1020.

ตัวอย่างที่ 5 การอ้างอิงเว็บไซต์

[13] B. Suttamanatwong. "Minimum Wage Rate 2022." (in Thai), MOL.go.th. <https://www.mol.go.th/อัตราค่าจ้างขั้นต่ำ> (accessed Dec. 31, 2022).

[14] Greenpeace Thailand. "The different air quality measurement standards." (in Thai), GREENPEACE.org. <https://www.greenpeace.org/thailandexplore/protect/cleanair/air-standard/> (accessed Nov. 23, 2022).

6. วิธีการจัดส่งบทความ

ผู้เขียนส่งบทความออนไลน์ได้ที่ <https://ph01.tci-thaijo.org/index.php/TNIJournal>

ทุกบทความที่ส่งเข้ามาจะมีการตรวจพิจารณาเบื้องต้นโดยกองบรรณาธิการก่อน เมื่อบทความผ่านการพิจารณาเบื้องต้นผู้เขียนจึงจะสามารถชำระเงินได้ หลังจากที่คุณชำระเงินเรียบร้อยแล้ว ทางกองบรรณาธิการจะดำเนินการส่งบทความเสนอผู้ทรงคุณวุฒิพิจารณาบทความและแจ้งผลการพิจารณาให้ผู้เขียนบทความทราบ สำหรับบทความที่ผ่านการประเมินโดยผู้ทรงคุณวุฒิแล้วจะได้รับการตีพิมพ์ลงในวารสารเพื่อเผยแพร่ต่อไป

วารสารวิศวกรรมศาสตร์และเทคโนโลยีดิจิทัล

ปีที่ 13 ฉบับที่ 2 กรกฎาคม – ธันวาคม 2568

“สร้างนักคิด ผลิตนักปฏิบัติ พัฒนานักบริหาร”

คณะวิศวกรรมศาสตร์

หลักสูตรระดับปริญญาตรี

- หลักสูตรวิศวกรรมยานยนต์ (Automotive Engineering, B.Eng.: AE)
- หลักสูตรวิศวกรรมหุ่นยนต์และระบบอัตโนมัติแบบสลับ (Robotics and Lean Automation Engineering, B.Eng.: RE)
- หลักสูตรวิศวกรรมคอมพิวเตอร์และปัญญาประดิษฐ์ (Computer Engineering and Artificial Intelligence, B.Eng.: CE)
- หลักสูตรวิศวกรรมอุตสาหการ (Industrial Engineering, B.Eng.: IE)
- หลักสูตรวิศวกรรมไฟฟ้า (Electrical Engineering, B.Eng.: EE)
- หลักสูตรวิศวกรรมบูรณาการระบบดิจิทัล (Digital System Integration Engineering, B.Eng.: SE)

หลักสูตรระดับปริญญาโท

- หลักสูตรเทคโนโลยีวิศวกรรม (Engineering Technology, M.Eng.: MET)

คณะเทคโนโลยีสารสนเทศ

หลักสูตรระดับปริญญาตรี

- หลักสูตรเทคโนโลยีสารสนเทศ (Information Technology, B.Sc.: IT)
- หลักสูตรเทคโนโลยีมัลติมีเดีย (Multimedia Technology, B.Sc.: MT)
- หลักสูตรเทคโนโลยีสารสนเทศทางธุรกิจดิจิทัล (Digital Business Information Technology, B.Sc.: BI)
- หลักสูตรเทคโนโลยีดิจิทัลทางสื่อสารมวลชน (Digital Technology in Mass Communication, B.Sc.: DC)
- หลักสูตรวิทยาการข้อมูลและการวิเคราะห์ข้อมูล (Data Science and Data Analytics, B.Sc.: DS)

หลักสูตรระดับปริญญาโท

- หลักสูตรเทคโนโลยีสารสนเทศ (Information Technology, M.Sc.: MIT)

คณะบริหารธุรกิจ

หลักสูตรระดับปริญญาตรี

- หลักสูตรบริหารธุรกิจญี่ปุ่น (Japanese Business Administration, B.B.A.: BJ)
- หลักสูตรการจัดการธุรกิจระหว่างประเทศ (International Business Management, B.B.A.: IB)
- หลักสูตรการบัญชี (Accountancy, B.Acc.: AC)
- หลักสูตรการจัดการทรัพยากรมนุษย์แบบญี่ปุ่น (Japanese Human Resources Management, B.B.A.: HR)
- หลักสูตรการตลาดดิจิทัล (Digital Marketing, B.B.A.: DM)
- หลักสูตรการจัดการการท่องเที่ยวและการบริการเชิงนวัตกรรม (Innovative Tourism and Hospitality Management, B.B.A.: TH)
- หลักสูตรการพัฒนาธุรกิจและสตาร์ทอัพ (Development of Business and Start-Up, B.B.A.: DBS)

หลักสูตรระดับปริญญาโท

- หลักสูตรบริหารธุรกิจญี่ปุ่น (Japanese Business Administration, M.B.A.: MBJ)
- หลักสูตรนวัตกรรมการจัดการธุรกิจและอุตสาหกรรม (Innovation of Business and Industrial Management, M.B.A.: MBI)
- หลักสูตรการจัดการระบบการผลิตและโลจิสติกส์แบบสลับ (Lean Manufacturing System and Logistics Management, M.B.A.: LMS)

คณะสื่อสารสากล

หลักสูตรระดับปริญญาตรี

- หลักสูตรภาษาญี่ปุ่นเพื่อธุรกิจระหว่างประเทศ (Japanese for International Business, B.A.: JIB)

คณะเทคโนโลยีดิจิทัลเพื่อธุรกิจอุตสาหกรรม

หลักสูตรระดับปริญญาตรี

- หลักสูตรโลจิสติกส์และดิจิทัลซัพพลายเชน (Logistics and Digital Supply Chain, B.S.: LD)

Thai-Nichi International College (TNIC)

Bachelor's Degree Programs

- สาขาวิศวกรรมดิจิทัล (Digital Engineering, B.Eng.: DGE)
- สาขาวิทยาการข้อมูลและปัญญาประดิษฐ์ (Data Science and Artificial Intelligence, B.Sc.: DSA)
- สาขาธุรกิจระหว่างประเทศและการเป็นผู้ประกอบการ (International Business and Entrepreneurship, B.B.A.: IBN)

Thai-Nichi Institute of Technology

1771/1 Pattanakarn Road, Suanluang, Bangkok 10250, Thailand

Tel: 0-2763-2600 Fax: 0-2763-2700

Website: <https://ph01.tci-thaijo.org/index.php/TNIJournal> E-mail : JEDT@tni.ac.th