



ISSN 2985-1580 (Print) ISSN 2985-1599 (Online)

วารสารคอมพิวเตอร์และเทคโนโลยีสร้างสรรค์

Journal of Computer and Creative Technology

วัตถุประสงค์วารสาร เพื่อส่งเสริมและสนับสนุนให้คณาจารย์ นักวิชาการ นิสิต นักศึกษา และผู้สนใจทั่วไป
ได้มีโอกาสเผยแพร่ผลงานทางวิชาการ

ขอบเขตวารสาร ดังนี้

- ❖ การประยุกต์ใช้คอมพิวเตอร์และเทคโนโลยีเพื่อการวิจัย พัฒนา และสร้างสรรค์
- ❖ การศึกษาและบูรณาการสู่การเรียนรู้ตลอดชีวิต
- ❖ สหวิทยาการเพื่อการพัฒนาท้องถิ่นและสังคม

กำหนดการออกวารสาร

ออกปีละ 3 ฉบับ ดังนี้ ฉบับที่ 1 ประจำเดือน มกราคม ถึง เมษายน (Issue 1— January to April)
ฉบับที่ 2 ประจำเดือน พฤษภาคม ถึง สิงหาคม (Issue 2— May to August)
ฉบับที่ 3 ประจำเดือน กันยายน ถึง ธันวาคม (Issue 3—September to December)

ประเภทบทความ

บทความวิจัย (Research Articles) และ บทความวิชาการ (Academic Articles)

การรับตีพิมพ์บทความ

บทความภาษาไทย (Thai Articles) และ บทความภาษาอังกฤษ (English Articles)

เงื่อนไขตีพิมพ์บทความ

บทความที่ส่งเข้าจะได้รับการพิจารณาเบื้องต้นโดยกองบรรณาธิการเกี่ยวกับขอบเขตของวารสารและรูปแบบการเขียนบทความ จากนั้นบทความที่ผ่านการพิจารณาเบื้องต้นจะถูกส่งไปพิจารณาประเมินคุณภาพของบทความจากคณะกรรมการผู้ทรงคุณวุฒิ (Peer Reviewer) อย่างน้อยสามท่านที่มีความเชี่ยวชาญตรงตามสาขาวิชาที่เกี่ยวข้องจากหลากหลายสถาบัน รูปแบบการประเมินบทความเป็นแบบผู้ทรงคุณวุฒิไม่ทราบชื่อผู้นิพนธ์และผู้นิพนธ์ไม่ทราบชื่อผู้ทรงคุณวุฒิ (Double-Blind Peer Review)

กองบรรณาธิการขอสงวนสิทธิ์ที่จะไม่พิจารณาประเมินบทความที่เคยตีพิมพ์เผยแพร่ในวารสารหรือสิ่งพิมพ์อื่น หรือกำลังอยู่ในระหว่างการพิจารณาเสนอขอตีพิมพ์ในวารสารหรือสิ่งพิมพ์อื่น ๆ นอกจากนี้ วารสารขอให้ผู้นิพนธ์เคร่งครัดในการเขียนบทความภายใต้สัญญาอนุญาตครีเอทีฟคอมมอนส์ (Attribution-NonCommercial-NoDerivatives 4.0 International : CC BY-NC-ND 4.0) หากมีการใช้ข้อมูลของผู้นิพนธ์อื่นจะต้องระบุที่มา โดยห้ามดัดแปลง ห้ามใช้เพื่อการค้า พร้อมทั้งขออนุญาตเจ้าของลิขสิทธิ์และแสดงหนังสือที่ได้รับการยินยอมต่อกองบรรณาธิการก่อนที่บทความจะได้รับการตีพิมพ์

ทั้งนี้ วารสารคอมพิวเตอร์และเทคโนโลยีสร้างสรรค์ ไม่มีค่าธรรมเนียมในการตีพิมพ์

คณะกรรมการวารสาร

ที่ปรึกษา	รองศาสตราจารย์ ดร.ฉลอง สุขทอง	อธิการบดีมหาวิทยาลัยราชภัฏสุรินทร์
	รองศาสตราจารย์ ดร.จิรายุ ทรัพย์สิน	รองอธิการบดีฝ่ายยุทธศาสตร์การยกระดับคุณภาพการศึกษา
	รองศาสตราจารย์ (พิเศษ) ดร.พรชัย เจตมานัน	นักวิชาการอิสระ อาจารย์พิเศษ ที่ปรึกษาสถาบันวิจัยและพัฒนา มหาวิทยาลัยราชภัฏมหาสารคาม และผู้ตรวจทานทีมสกอปัส
	ผู้ช่วยศาสตราจารย์ประยุทธิ์ คงอินทร์	คณบดีคณะวิทยาศาสตร์และเทคโนโลยี
	รองคณบดีคณะวิทยาศาสตร์และเทคโนโลยี	

บรรณาธิการ	อาจารย์ ดร.วิจิตรา โพธิสาร	มหาวิทยาลัยราชภัฏสุรินทร์
ผู้ช่วยบรรณาธิการ	ผู้ช่วยศาสตราจารย์ ดร.วาฤทธิ์ นวลนาง	มหาวิทยาลัยราชภัฏสุรินทร์
	ผู้ช่วยศาสตราจารย์ ดร.ธงชัย เจือจันทร์	มหาวิทยาลัยราชภัฏสุรินทร์

กองบรรณาธิการ

ภายนอกมหาวิทยาลัย

รองศาสตราจารย์ ดร.สุรพล บุญลือ	มหาวิทยาลัยเทคโนโลยีพระจอมเกล้าธนบุรี
รองศาสตราจารย์ ดร.สิทธิชัย บุษหมั่น	มหาวิทยาลัยราชภัฏมหาสารคาม
รองศาสตราจารย์ ดร.เนติรัฐ วีระนาคินทร์	มหาวิทยาลัยมหาสารคาม
ผู้ช่วยศาสตราจารย์ ดร.สมนึก พ่วงพรพิทักษ์	มหาวิทยาลัยมหาสารคาม
ผู้ช่วยศาสตราจารย์ ดร.สุวิธ ธีระโคตร	มหาวิทยาลัยมหาสารคาม
ผู้ช่วยศาสตราจารย์ ดร.สุวัฒน์ บรรลือ	มหาวิทยาลัยราชภัฏอุบลราชธานี
ผู้ช่วยศาสตราจารย์ ดร.เกษศิริ ทองเฉลิม	มหาวิทยาลัยราชภัฏอุบลราชธานี
ผู้ช่วยศาสตราจารย์ ดร.อนล สอนประดิษฐ์	มหาวิทยาลัยราชภัฏบุรีรัมย์
ผู้ช่วยศาสตราจารย์ ดร.ประวิทย์ สิมมาทัน	มหาวิทยาลัยราชภัฏมหาสารคาม
อาจารย์ ดร.อัจฉรา ชุมพล	มหาวิทยาลัยกาฬสินธุ์
อาจารย์ ดร.ปิยวัฒน์ คำสบาย	มหาวิทยาลัยราชภัฏอุดรธานี
อาจารย์ ดร.วิภาสทธิ์ หิรัญรัตน์	มหาวิทยาลัยเทคโนโลยีราชมงคลอีสาน วิทยาเขตสุรินทร์
อาจารย์นวัตร โพธิสาร	มหาวิทยาลัยเทคโนโลยีราชมงคลอีสาน วิทยาเขตสุรินทร์

ภายในมหาวิทยาลัย

ผู้ช่วยศาสตราจารย์ ดร.ขจรศักดิ์ สงวนสัตย์	มหาวิทยาลัยราชภัฏสุรินทร์
ผู้ช่วยศาสตราจารย์ ดร.วันชัย สุขตาม	มหาวิทยาลัยราชภัฏสุรินทร์
ผู้ช่วยศาสตราจารย์ ดร.ศิริลักษณ์ หวังขอบ	มหาวิทยาลัยราชภัฏสุรินทร์
อาจารย์ ดร.ประภัสญา สร้อยจิตร	มหาวิทยาลัยราชภัฏสุรินทร์
อาจารย์ ดร.จักรพงษ์ วารี	มหาวิทยาลัยราชภัฏสุรินทร์

ทีมสนับสนุนงานวารสาร

ประสานงานบทความ : อาจารย์กัญญาณี สมอ
สื่อดิจิทัล/ เว็บไซต์วารสาร : อาจารย์กฤษฎา นวลนาง
ประชาสัมพันธ์ : อาจารย์สุวัฒน์ กล้ายทอง
เลขานุการ : นายทวิวัฒน์ มูลจืด

สำนักงาน

สาขาวิชาคอมพิวเตอร์ศึกษา
คณะวิทยาศาสตร์และเทคโนโลยี มหาวิทยาลัยราชภัฏสุรินทร์
186 หมู่ 1 ถ.สุรินทร์-ปราสาท ต.นอกเมือง อ.เมือง จ.สุรินทร์ 32000
โทรศัพท์ 044 – 558344, 088 723 5944 อีเมล jcct@srru.ac.th



บทบรรณาธิการ

วารสารคอมพิวเตอร์และเทคโนโลยีสร้างสรรค์ (Journal of Computer and Creative Technology) เปิดรับบทความในขอบเขต ดังนี้ การประยุกต์ใช้คอมพิวเตอร์และเทคโนโลยีเพื่อการวิจัย พัฒนา และสร้างสรรค์ การศึกษาและบูรณาการสู่การเรียนรู้ตลอดชีวิต และสหวิทยาการเพื่อการพัฒนาท้องถิ่นและสังคม โดยกำหนดออกปีละ 3 ฉบับ ได้แก่ ฉบับที่ 1 ประจำเดือน มกราคม - เมษายน ฉบับที่ 2 ประจำเดือน พฤษภาคม - สิงหาคม และฉบับที่ 3 ประจำเดือน กันยายน - ธันวาคม ทั้งบทความวิจัยและบทความวิชาการ รูปแบบการตีพิมพ์บทความภาษาไทยและบทความภาษาอังกฤษ สำหรับฉบับนี้เป็นปีที่ 2 ฉบับที่ 2 ประจำเดือน พฤษภาคม - สิงหาคม พ.ศ. 2567 มีบทความ จำนวน 5 บทความ เป็นบทความวิจัย จำนวน 5 เรื่อง ได้แก่ 1) เครื่องจ่ายยาอัตโนมัติและตรวจสอบความถูกต้องด้วยโมเดล YOLOv5 2) Comparing a Thai Words Segmentation Methods in the LST20 Dataset 3) การทำนายราคาที่ดินด้วยการเรียนรู้ของเครื่องในอำเภอเมืองขอนแก่น 4) การพัฒนาผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ ความสามารถในการแก้ปัญหาทางคณิตศาสตร์ และความรับผิดชอบโดยใช้การจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1 และ 5) การพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำรินหลวง

กองบรรณาธิการหวังเป็นอย่างยิ่งว่าบทความที่เผยแพร่ในวารสารคอมพิวเตอร์และเทคโนโลยีสร้างสรรค์ จะเป็นแหล่งเรียนรู้ออนไลน์ที่สำคัญทางคอมพิวเตอร์และเทคโนโลยีสร้างสรรค์ เป็นแหล่งเผยแพร่ทางวิชาการให้กับนักวิจัยและนักวิชาการ และเป็นประโยชน์ให้กับผู้อ่านและการนำไปใช้ประโยชน์ต่อไป

อาจารย์ ดร.วิจิตรา โพธิสาร
บรรณาธิการ

สารบัญ

	หน้า
บทนำ	I-V
บทความวิจัย	
● เครื่องจ่ายยาอัตโนมัติและตรวจสอบความถูกต้องด้วยโมเดล YOLOv5 ปิ่นพงษ์ เรืองระวีณุกิจ, ธนกฤต นุพิมพ์, อนุชา ลือบางใหญ่, พลวัต ช่อผุก	45-60
● Comparing a Thai Words Segmentation Methods in the LST20 Dataset Krittapol Damrongkamoltip, Khatcha Ruenlek, Wasit Limprasert, Prachya Boonkwan	61-70
● การทำนายราคาที่ดินด้วยการเรียนรู้ของเครื่องในอำเภอเมืองขอนแก่น โยษิตา ศรีวุฒิทรัพย์, ดุสิตา สังข์กลิ่นหอม, ศักดิ์พจน์ ทองเลี่ยมนาค, ธนพล ตั้งชูพงศ์	71-86
● การพัฒนาผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ ความสามารถในการแก้ปัญหาทาง คณิตศาสตร์ และความรับผิดชอบโดยใช้การจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียน ชั้นมัธยมศึกษาปีที่ 1 ชฎานิศ เมืองคู่, ยุภาดี ปณะราช	87-98
● การพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง บุญทฤก สีหาบุตรโต, กมลรัตน์ สมใจ, นิธินันท์ มาตา	99-108



Citation:

บทความในวารสารนี้อยู่ภายใต้สัญญาอนุญาตครีเอทีฟคอมมอนส์ (Attribution-NonCommercial-NoDerivatives 4.0 International : CC BY-NC-ND 4.0)

Automated Medicine Dispenser and Verification by YOLOv5 Models

เครื่องจ่ายยาอัตโนมัติและตรวจสอบความถูกต้องด้วยโมเดล YOLOv5

Pinphong Ruangraweenukit¹, Thanakrit Nupim², Anucha Luebangyai³, and Ponlawat Chophuk^{4,*}

ปิ่นพงษ์ เรืองระวีณุกิจ¹, ธนกรุต นูพิมพ์², อนุชา ลือบางไญ้³, และ พลวัต ช่อผุก^{4,*}

Received: 9 February 2024;

Revised: 21 April 2024;

Accepted: 25 April 2024;

Published: 18 July 2024;

Abstract

Medication errors are a major problem to affects the safety of patients around the world, and is one of the causes of death is incorrect drug use. This issue is particularly acute for patients with complex medication schedules, exacter Medication errors are a major problem to affects the safety of patients around the world, and is one of the causes of death is incorrect drug use bated by hospital dispensing delays that contribute to these errors. This research focuses on developing an automated medicine dispenser powered by a Raspberry Pi to enhance the accuracy of drug dispensing, reduce errors, and provide pertinent medication information. The system is designed to accept direct prescription inputs from physicians and incorporates YOLOv5 as the artificial intelligence (AI) component to ensure prescription accuracy. The findings from the development phase reveal that Model 3 demonstrated the highest efficiency at 64%. The AI's verification accuracy was determined to be 90%. Which makes the medication error from the system 11%.

Keywords: Automated Medicine, Dispenser and Verification, YOLOv5 Models.

บทคัดย่อ

ความคลาดเคลื่อนทางยา คือ ปัญหาสำคัญที่ส่งผลกระทบต่อความปลอดภัยของผู้ป่วยทั่วโลก และเป็นสาเหตุหนึ่งของการเสียชีวิตจากการใช้ยาไม่ถูกต้อง ปัญหานี้ยังเพิ่มความซับซ้อนสำหรับผู้ที่ต้องการการรักษาด้วยสูตรยาที่มีความซับซ้อน การวิจัยนี้มุ่งเน้นที่

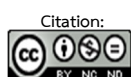
¹ Student, Faculty of Informatics, Burapha University, Chonburi 20131, Thailand; นักศึกษา คณะวิทยาการสารสนเทศ มหาวิทยาลัยบูรพา จังหวัดชลบุรี 20131 ประเทศไทย; Email: 64160205@go.buu.ac.th

² Student, Faculty of Informatics, Burapha University, Chonburi 20131, Thailand; นักศึกษา คณะวิทยาการสารสนเทศ มหาวิทยาลัยบูรพา จังหวัดชลบุรี 20131 ประเทศไทย; Email: 64160118@go.buu.ac.th

³ Student, Faculty of Informatics, Burapha University, Chonburi 20131, Thailand; นักศึกษา คณะวิทยาการสารสนเทศ มหาวิทยาลัยบูรพา จังหวัดชลบุรี 20131 ประเทศไทย; Email: 64160211@go.buu.ac.th

^{4,*} Lecturer, Faculty of Informatics, Burapha University, Chonburi 20131, Thailand; อาจารย์ คณะวิทยาการสารสนเทศ มหาวิทยาลัยบูรพา จังหวัดชลบุรี 20131 ประเทศไทย; Email: ponlawat.ch@go.buu.ac.th

*Corresponding authors: Ponlawat Chophuk (ponlawat.ch@go.buu.ac.th)



การพัฒนาเครื่องจ่ายยาอัตโนมัติที่ใช้รหัสบาร์โค้ดเพื่อปรับปรุงประสิทธิภาพในการจ่ายยา ลดความผิดพลาดโดยการใช้ปัญญาประดิษฐ์ ระบบที่พัฒนาสามารถรับรายการคำสั่งจ่ายยาจากแพทย์โดยตรง และใช้ YOLOv5 เป็นเครื่องมือปัญญาประดิษฐ์ (AI) ในการตรวจสอบความถูกต้องของรายการยา ผลลัพธ์จากการพัฒนาโมเดลแสดงให้เห็นว่า โมเดลที่สามมีประสิทธิภาพสูงสุดที่ 64% การตรวจสอบความถูกต้องโดยใช้ AI มีความแม่นยำถึง 90%ซึ่งทำให้มีความคลาดเคลื่อนทางยาอยู่ที่ 11%

คำสำคัญ: การจ่ายยาอัตโนมัติ, เครื่องจ่ายยาและการตรวจสอบ, โมเดล YOLOv5

1. บทนำ (Introduction)

ความคลาดเคลื่อนทางยาเป็นปัญหาสากลที่มีผลกระทบต่อสุขภาพและความเป็นอยู่ของผู้ป่วย ปัญหานี้ส่งผลให้เกิดความจำเป็นในการดำเนินการอย่างจริงจังทั้งในด้านสาธารณสุขและความปลอดภัยของผู้ป่วย ความคลาดเคลื่อนทางยาสามารถเกิดขึ้นได้จากหลายปัจจัย ตั้งแต่ความผิดพลาดในการจ่ายยา ความล้มเหลวในการสื่อสารระหว่างบุคลากรทางการแพทย์ ไปจนถึงความเหนื่อยล้าและความเครียด และความผิดพลาดในการใช้ยาโดยผู้ป่วยเอง ดังนั้นเพื่อให้ปัญหาและวัตถุประสงค์ของงานวิจัยชัดเจนยิ่งขึ้น ขอชี้แจงว่าปัญหาความคลาดเคลื่อนทางยานี้ได้รับการยอมรับและนำเสนอเป็นประเด็นวิจัยในหลายหน่วยงานสาธารณสุขและสถาบันการศึกษาทั่วโลก การศึกษาย้อนหลังจากองค์การเพื่อความร่วมมือทางเศรษฐกิจและการพัฒนา (OECD) ในปี 2022 ชี้ให้เห็นว่าความคลาดเคลื่อนทางยาเป็นปัญหาที่ต้องได้รับการแก้ไขอย่างเร่งด่วนเพื่อลดความสูญเสียที่อาจเกิดขึ้น ซึ่งประเมินได้ถึง 54 พันล้านดอลลาร์สหรัฐ (de Bienassis et al., 2022)

การสนับสนุนการใช้เทคโนโลยีใหม่ในการแก้ไขปัญหาคความคลาดเคลื่อนทางยา มีการศึกษาหลายชิ้นที่ได้นำเสนอการใช้ Deep Learning ในด้านการแพทย์และยา ตัวอย่างเช่น การวิจัยที่ตีพิมพ์ใน สถาบันวิชาชีพวิศวกรไฟฟ้าและอิเล็กทรอนิกส์ (Institute of Electrical and Electronics Engineers: IEEE) ได้แสดงให้เห็นว่าระบบที่พัฒนาด้วย Deep Learning สามารถช่วยในการตรวจสอบและระบุความผิดปกติของร่างกายได้ (Ravi et al., 2017) อย่างมีประสิทธิภาพ ด้วยการฝึกอบรมระบบด้วยข้อมูลจำนวนมาก ระบบเหล่านี้สามารถเรียนรู้และปรับปรุงความแม่นยำในการตรวจจับความคลาดเคลื่อนได้ ซึ่งจะนำไปสู่การลดข้อผิดพลาดและเพิ่มความปลอดภัยให้กับผู้ป่วย

นอกจากนี้ การวิเคราะห์เบื้องต้นและการทบทวนวรรณกรรมเกี่ยวกับเครื่องจ่ายยาอัตโนมัติได้ชี้ให้เห็นว่าแม้จะมีการใช้งานอยู่แล้วในบางสถานที่ แต่ยังมีข้อจำกัดและปัญหาในเรื่องของความแม่นยำและความสามารถในการตรวจสอบรายการยาอย่างละเอียด ดังนั้น การพัฒนาเครื่องจ่ายยาที่ปรับปรุงด้วยเทคโนโลยีปัญญาประดิษฐ์จึงเป็นวิธีการที่น่าสนใจในการแก้ไขปัญหานี้

2. วัตถุประสงค์งานวิจัย (Research Objectives)

1. เพื่อสร้างเครื่องจ่ายยาอัตโนมัติและตรวจสอบโดยการใช้ปัญญาประดิษฐ์
2. เพื่อลดปัญหาความคลาดเคลื่อนทางยาที่เกิดจากมนุษย์ เช่น ความเหนื่อยล้าจากการปฏิบัติงาน ภาระทางร่างกายหรืออารมณ์ของบุคลากรทางการแพทย์ และการสื่อสารระหว่างบุคลากรทางการแพทย์เป็นต้น

3. กรอบแนวคิดงานวิจัย (Conceptual Framework)

การจ่ายยาเป็นสิ่งสำคัญหากมีการผิดพลาดอาจส่งผลกระทบต่อผู้ป่วยได้ และการขาดความรู้ทางยาในการรับประทานจะส่งผลเสียอันตรายต่อผู้ป่วยได้ ดังนั้นควรมีระบบในการเก็บข้อมูลตัวยา และในกระบวนการจ่ายยาควรมีขั้นตอนการตรวจสอบเพื่อป้องกันไม่ให้เกิดการจ่ายยาผิดพลาด ดังนี้

1. ระบบการสั่งยาผ่าน QR Code ช่วยลดเวลาในการเตรียมยาโดยอัตโนมัติเมื่อสแกน QR ด้วยมือถือ และสามารถใช้อินเตอร์เฟซเว็บไซต์ที่พัฒนาขึ้น
2. ระบบฐานข้อมูลยา สามารถเก็บรวบรวมข้อมูลตัวยาที่ใช้งานในเครื่องจ่ายยาและอำนวยความสะดวกในการเพิ่มยาตัวใหม่เข้าสู่ระบบ โดยใช้ระบบการจัดการฐานข้อมูล MySQL
3. ระบบจ่ายยาอัตโนมัตินี้ได้รับการออกแบบมาสำหรับยาเม็ดที่มีขนาดเล็กกว่า 2 เซนติเมตร โดยใช้โมดูลเซอร์โมเตอร์และ Raspberry Pi ในการควบคุมการทำงานของเครื่องจ่ายยา
4. สามารถใช้ปัญญาประดิษฐ์ในการตรวจสอบความถูกต้องของตัวยาได้ โดยใช้โมเดล YOLOv5s ซึ่งเป็นการเรียนรู้เชิงลึกสำหรับการรับรู้ภาพ และนำระบบ Raspberry Pi และโมดูลกล้อง มาใช้เป็นฮาร์ดแวร์หลักในการดำเนินการตรวจจับ

4. การทบทวนวรรณกรรมและทฤษฎีที่เกี่ยวข้อง (Literature Review)

ในยุคที่เทคโนโลยีเติบโตอย่างไม่หยุดยั้ง การพัฒนาระบบการสั่งยาออนไลน์กลายเป็นหนึ่งในด้านที่ได้รับความสนใจอย่างมากจากผู้วิจัยหลายคน เพื่อเพิ่มประสิทธิภาพและความถูกต้องในการจ่ายยาให้กับผู้ป่วย ซึ่งเริ่มต้นจากการวิจัยของ Niu et al. (2023) ที่แสดงให้เห็นถึงระบบการสั่งยาออนไลน์ที่มีการตรวจสอบซ้ำ เพื่อให้ยาที่จ่ายออกไปตรงกับสั่งของแพทย์ แม้ว่าระบบนี้จะช่วยให้ผลลัพธ์ออกมาตรงกับที่ผู้ป่วยต้องการ แต่ก็พบว่ามีข้อจำกัดคือใช้เวลานานเกินไปในการทำงาน ในการแก้ไขปัญหาดังกล่าว Zheng et al. (2023) ได้ทดลองใช้ปัญญาประดิษฐ์ โดยให้มนุษย์เป็นศูนย์กลางของระบบ เพื่อป้องกันการจ่ายยาผิด โดยมีเมสซิงเจอร์คอยกำกับการดูแล และใช้ Computer Vision เพื่อตรวจจับยา แต่ก็ยังพบปัญหาเนื่องจากยาที่มีความคล้ายคลึงกันทำให้เกิดความผิดพลาดในการตรวจจับ ต่อมา Bu et al. (2022) ได้พัฒนาระบบจ่ายยาผ่านอินเทอร์เน็ตด้วยการใช้ปัญญาประดิษฐ์ ทดสอบประสิทธิภาพในโรงพยาบาลในจีน โดยเฉพาะการยืนยันปริมาณการใช้ยาและการออกใบสั่งยา อย่างไรก็ตาม ปัญหาที่พบคือผู้สูงอายุมีความยากลำบากในการใช้งาน และต้องมีการตรวจสอบการทำงานของ AI อย่างต่อเนื่อง ซึ่ง Kim et al. (2022) ได้พัฒนาเป็นอีกขั้นด้วยการใช้การยืนยันด้วยใบหน้าในเครื่องจ่ายยาเพื่อเพิ่มความแม่นยำและการแจ้งเตือนแบบ Real-time ผ่านแอปพลิเคชัน แม้จะมีความซับซ้อนในการเริ่มต้นใช้งานระบบตรวจจับใบหน้า นอกจากนี้ Nasir et al. (2023) นำเสนอ Smart Medical Box ที่ใช้ระบบการจำแนกด้วยใบหน้าเพื่อยืนยันตัวตนผู้ป่วย และส่ง SMS เมื่อมียาถูกจ่ายออกไป อย่างไรก็ตาม การใช้งานฟังก์ชัน Biometric Recognition ยังคงซับซ้อนและผู้ป่วยมีความเข้าใจน้อย และ Dayananda & Upadhya (2024) ได้พัฒนาระบบ Smart Pill ด้วยการใช้ IoT เพื่อแก้ปัญหาที่เกิดขึ้นก่อนหน้านี้ โดยมุ่งเน้นที่การใช้งานที่เข้าใจง่ายสำหรับทุกวัย และการป้องกันการ Overdose และการดื้อยา แต่ยังมีข้อจำกัดในการ Verify ผลการตรวจจับยา ท้ายที่สุด Anupama et al. (2020) และ Gargioni et al. (2024) ได้ทำการวิเคราะห์รวมถึงประโยชน์ของ AI ในการแพทย์ สรุปได้ว่า AI เพิ่มประสิทธิภาพและความแม่นยำในการตัดสินใจลดงบประมาณในการดูแลระบบ และทำนายโรคได้ ซึ่งจุดที่ต้องการในตลาดคือการมีระบบการใช้งานที่เข้าใจง่ายและคำแนะนำการใช้ยา ประกอบกับงานวิจัยของ Mumford et al. (2024) ที่ได้อภิปรายเกี่ยวกับขั้นตอนการวัดความคลาดเคลื่อนทางยา พบว่า จาก 26,369 กรณี มีถึง 19,692 กรณีที่มีความคลาดเคลื่อน คิดเป็น 74% และมี 1,473 กรณี (7.5%) ที่มีผลข้างเคียงร้ายแรง จากนั้น Barker et al. (1984) จึงได้ทำการทดสอบวัดประสิทธิภาพเครื่องให้ยาอัตโนมัติของแพทย์พบว่ามีความคลาดเคลื่อนทางยาอยู่ที่ 10.6% และส่วนใหญ่เป็นความคลาดเคลื่อนที่เกิดจากเวลา

จากภาพรวมของการวิจัยเหล่านี้ งานวิจัยจึงได้พัฒนาเครื่องจ่ายยาที่สามารถส่งงานได้ผ่าน QR code ซึ่งแพทย์เป็นผู้สั่งการ ผ่านเว็บไซต์ของเครื่องจ่ายยาโดยตรง ที่นอกจากจะมีข้อมูลยาต่าง ๆ ยังช่วยให้ผู้ป่วยสามารถติดตามอาการและวินิจฉัยโรคได้ง่ายขึ้น สะท้อนถึงการวิวัฒนาการของเทคโนโลยีในด้านการแพทย์ที่มุ่งเน้นทั้งในเรื่องของความแม่นยำ ความง่ายในการใช้งาน และการดูแลผู้ป่วยอย่างใกล้ชิด

4.1 YOLOv5

YOLOv5 (โยโลเวอร์ชัน 5) เป็น Algorithms ที่ถูกสร้างโดย Glenn-Jocher ด้วยการประมวลผลแบบ Real-time ในด้านการทำงานของ Algorithms (Chen et al., 2022; Xu et al, 2022) YOLOv5 เป็นส่วนประกอบที่สำคัญในการทำงานของเครื่องจ่ายยาอัตโนมัติ เนื่องจากการใช้การตรวจจบบรรยากาศ ทั้งในการฝึก (Training) และการวัดผล (Detection) YOLOv5 ประกอบด้วย 4 ส่วนหลัก คือ

1. Input Side ทำหน้าที่ในการรับข้อมูลโดยเป็น จะรับเป็นไฟล์รูปภาพ
2. Backbone Network ทำหน้าที่ในการสกัดคุณสมบัติจากภาพที่ Input เข้ามา
3. Neck Network ทำหน้าที่ปรับปรุงคุณสมบัติที่ได้มาจาก Backbone เพื่อเพิ่มประสิทธิภาพให้มากขึ้น
4. Output Part ทำหน้าที่ประมวลผลคุณสมบัติจาก Neck เพื่อสร้างการคาดการณ์ (Prediction) (Chen et al., 2022)

4.2 Raspberry Pi

Raspberry Pi (ราสเบอร์รี่พาย) คือ คอมพิวเตอร์บอร์ดเดี่ยว สร้างสรรค์โดยมูลนิธิ Raspberry Pi บอร์ด Raspberry Pi แต่ละตัวประกอบด้วยส่วนประกอบที่สำคัญ เช่น หน่วยประมวลผลกลาง (CPU), หน่วยความจำเข้าถึงโดยสุ่ม (RAM), พอร์ต USB, เอาต์พุต HDMI และพินอินพุต/ เอาต์พุตวัตถุประสงค์ทั่วไป (GPIO) โดยรองรับระบบปฏิบัติการต่าง ๆ ซึ่งส่วนใหญ่เป็นระบบปฏิบัติการ Linux (Watson, 2018) โดยบอร์ด Raspberry Pi ถูกเลือกมาใช้ในการเชื่อมต่ออินเทอร์เน็ตและต่อกับโมดูลต่าง ๆ เพื่อนำมาเป็นเครื่องที่จะปล่อยยา ซึ่งในที่นี้คือ Pi Camera และ Servo

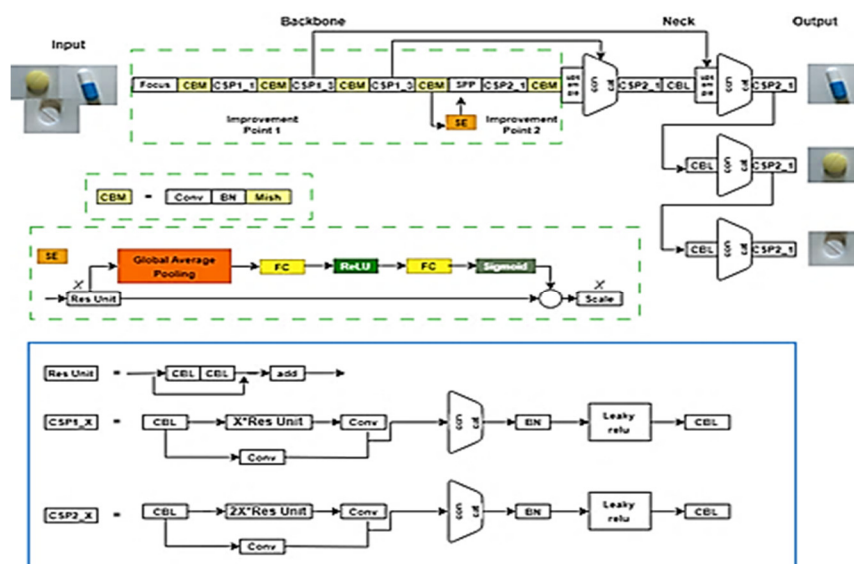


Figure 1. YOLOv5 architecture and workflow of Automated Medicine Dispenser.

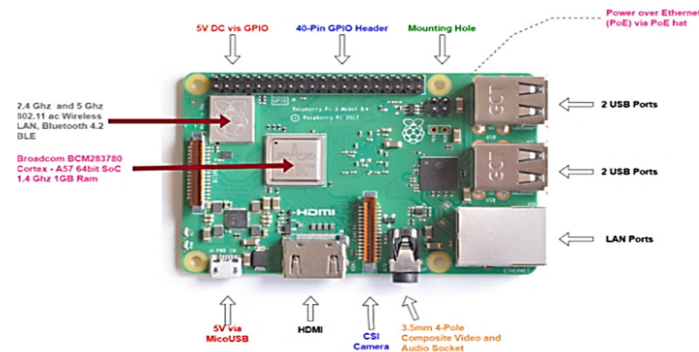


Figure 2. Raspberry Pi 3 Model B.

4.3 Database

Database คือการเก็บรวบรวมข้อมูลที่มีโครงสร้างซึ่งออกแบบมาเพื่อการจัดเก็บ และการจัดการที่มีประสิทธิภาพ โดยทำหน้าที่เป็นองค์ประกอบพื้นฐานในการจัดเก็บข้อมูล

Relational Database จะจัดการข้อมูลในตารางโดยใช้ SQL ในการเข้าถึงข้อมูล การจัดการข้อมูลที่เชื่อมโยงและกระจัดกระจายกันเช่น ข้อมูลยา ข้อมูลคนไข้ข้อมูลการออกไปส่งยา เป็นต้น (Anderson & Nicholson, 2022)

4.4 Rest API

Application Programming Interface หรือ API เป็นหนึ่งในวิธีในการติดต่อสื่อสารกันของ Software และจะมี Protocol ในการติดต่อสื่อสารชัดเจน โดย Application ที่ส่ง Request เรียกว่า Client และ Application ที่ส่ง Respond เรียกว่า Server ซึ่งประเภทของ API ที่นำมาใช้งานคือ REST API

REST ย่อมาจาก Representational State Transfer โดย REST นั้น จะมีฟังก์ชันหลาย อย่างเช่น GET, PUT, DELETE ที่ Client สามารถใช้เพื่อเข้าถึงข้อมูลเหล่านั้นได้ โดย Client และ Server จะเชื่อมต่อกันโดยใช้ HTTP (Amazon Web Services, n.d.) ฟังก์ชันหลักของ REST นั้น คือการไม่บันทึกข้อมูลฝั่ง Client ในการเรียกใช้แต่ละครั้ง การเรียกใช้ หรือ Request จากฝั่ง Client นั้น มักจะอยู่ในรูปแบบของ URL เพื่อให้ง่ายต่อการเข้าถึงข้อมูล ซึ่งเราสามารถนำมาประยุกต์ใช้ได้กับการทำ QR Code โดยให้ Request เป็น Link ที่ได้มาจาก QR Code เพื่อดึงข้อมูลรายละเอียดยา (Respond) ของใบสั่งในรูปแบบ JSON

4.5 Internet of things (IoT)

Internet of things (IoT) คือ การพัฒนาของเทคโนโลยีที่ขับเคลื่อนความก้าวหน้าทางด้านโทรคมนาคมและพัฒนาคุณภาพชีวิตของผู้คนทั่วโลก อีกทั้ง IoT นั้นยังศักยภาพในการฟื้นฟูเศรษฐกิจโลกอีกด้วย ซึ่งถูกคาดการณ์ไว้ว่ามีคุณค่าต่อเศรษฐกิจโลกถึง 2.7 ล้านล้าน USD และ 6.2 ล้านล้าน USD ในปี 2025 (Al-Fuqaha et al., 2015) เนื่องจากอุปกรณ์ IoT นั้น เป็นส่วนประกอบที่สำคัญของแอปพลิเคชันที่เกิดขึ้นใหม่ และมีบทบาทในการสื่อสารกันระหว่างเครื่องจักรไปใช้งานอย่างแพร่หลาย จากคำพูดของผู้เชี่ยวชาญ IoT นั้นมีศักยภาพมากพอที่จะปฏิวัติเทคโนโลยีอัจฉริยะต่าง ๆ เช่น Smart Home และ Modern Healthcare (Zanella et al., 2014) โดยเฉพาะ ในอุตสาหกรรมด้านสุขภาพ ซึ่งถูกคาดการณ์มูลค่าการตลาดไว้สูงถึง 2.5 ล้านล้าน USD ในปี 2025 (Almotiri et al., 2016)

4.6 Deep Learning

Deep learning คือการเรียนรู้แบบอัตโนมัติผ่านการเลียนแบบการทำงาน โดยการนำระบบโครงข่ายประสาท (Neural Network) มาซ้อนกันหลาย ๆ ชั้น (Layer) และทำการเรียนรู้จากข้อมูลตัวอย่างที่ถูกป้อนไปให้ ซึ่งข้อมูลดังกล่าวจะถูกนำไปใช้ในการตรวจจับรูปแบบ (Pattern) หรือจัดหมวดหมู่ข้อมูล (Classification) (ABB, 2020)

Deep learning จัดอยู่ในหมวดหมู่ Multilayer Neural Network ที่จะเน้นไปที่การเรียนรู้คุณสมบัติของข้อมูลที่ถูกส่งมาให้ (Shen et al., 2017) และได้มีการทดสอบการใช้ Deep learning ในทางการแพทย์หลายอย่าง เช่น ใช้ในการแบ่งประเภทของอวัยวะ, การตรวจจับโครงสร้างทางกายวิภาค, การตรวจจับเซลล์เป็นต้น (Shen et al., 2017)

5. วิธีดำเนินงานวิจัย (Research Methodology)

5.1 ศึกษาข้อมูลยากกลุ่มตัวอย่าง

งานวิจัยนี้ได้เลือก ยาชนิดผงรูปทรงกลมกับยาชนิดแคปซูลที่มีขนาดไม่เกิน 2 เซนติเมตร มาทำการศึกษาเกี่ยวกับชนิด ลักษณะ รูปทรง สรรพคุณของยา และได้ถ่ายรูปภาพยาด้วยกล้องถ่ายรูป เพื่อนำมาเก็บเป็นชุดข้อมูลยาที่ต้องการนำไปใช้ร่วมกับระบบ

5.2 เครื่องมือและหลักการที่ใช้

งานวิจัยนี้ได้มีการเลือกนำเครื่องมือและหลักการต่าง ๆ มาใช้ในการพัฒนาระบบเครื่องจ่ายยาอัตโนมัติ โดยที่จะแบ่งระบบออกเป็น 2 ส่วนดังนี้

1. ระบบสั่งยาผ่านเว็บไซต์ ได้มีการใช้เครื่องมือ เช่น ฐานข้อมูล Mysql ในการเก็บข้อมูลยา Rest API ในการติดต่อสื่อสารกับฐานข้อมูล Figma ในการออกแบบเว็บไซต์ และ VS Code ในการพัฒนา เป็นต้น

2. ระบบจ่ายยาอัตโนมัติ ได้มีการใช้เครื่องมือ เช่น Raspberry Pi ในการสั่งการจ่ายยาของตัวโมเดล ด้วยการควบคุม Servo Motor การสร้างเขียนแบบโมเดล 3 มิติด้วย Tinkercad วัสดุที่ใช้ในการทำโมเดลทั้ง ฟิวเจอร์บอร์ด กระดาษแข็ง และโฟม YOLOv5s สำหรับตรวจสอบความถูกต้องของยา เป็นต้น

5.3 การออกแบบระบบ

งานวิจัยนี้ได้มีการวางแผนออกแบบเกี่ยวกับการทำงานของระบบเครื่องจ่ายยาอัตโนมัติ โดยจะมีการออกแบบภาพหน้าเว็บไซต์หน้าต่าง ๆ ด้วย Figma การออกแบบโมเดล 3 มิติของเครื่องจ่ายยาด้วย Tinkercad การออกแบบสร้างแผนภาพไดอะแกรม ทั้งในส่วน of flowchart การทำงานของระบบด้วย Draw.io ตัวอย่างเช่นดัง Figure 3. ที่เกี่ยวกับภาพรวมของระบบ โดยจะแบ่งออกเป็น 2 ส่วน คือ 1. ระบบสั่งยาผ่านเว็บไซต์ และ 2. ระบบจ่ายยาอัตโนมัติผ่านโมเดลเครื่องจ่ายยา โดยทั้งสองส่วนจะมีการนำมาทำงานร่วมกัน โดยจะเห็นว่า ระบบหน้าเว็บไซต์สั่งยาจะออกไปสั่งยา โดยสร้าง QR Code ให้ในใบสั่งยา และระบบเครื่องจ่ายยาจะมีหน้าที่อ่าน QR Code จากใบสั่งยา แล้วเครื่องจะทำการจ่ายออกไป ผ่านตัวโมเดลที่เก็บยา จากภาพรวมของระบบข้างต้นที่ได้ทำการแบ่งออกเป็น 2 ส่วน จะถูกออกแบบเพิ่มเติมในแต่ละส่วนดังนี้

1. ระบบสั่งยาผ่านเว็บไซต์ ในระบบได้แบ่งออกมาเป็น 2 ส่วน โดยจะประกอบไปด้วย ระบบฐานข้อมูลจัดเก็บข้อมูลยา หมอ และคนไข้ ดัง Figure 4. จะเป็นการออกแบบ ER Diagram ด้วย Draw.io และระบบหน้าเว็บไซต์ ที่ประกอบไปด้วย เว็บไซต์หน้าต่าง ๆ ดัง Figure 5. ที่เป็นหน้าเว็บไซต์ที่ใช้ในการแสดงข้อมูลยา จะถูกออกแบบด้วย Figma ทั้งหมด ในหน้าเว็บไซต์ในการแสดงข้อมูลคนไข้ และหน้าเว็บไซต์ของหมอที่ใช้สำหรับออกไปสั่งยาให้คนไข้จะมีการออกแบบกระบวนการทำงาน ดัง Figure 6. ด้วย Draw.io

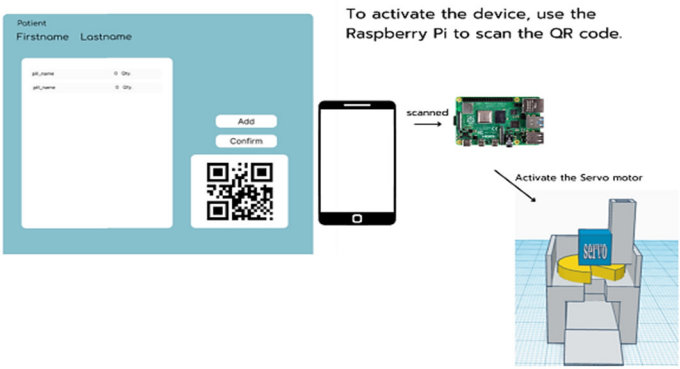


Figure 3. Automated Medicine Dispenser operation.

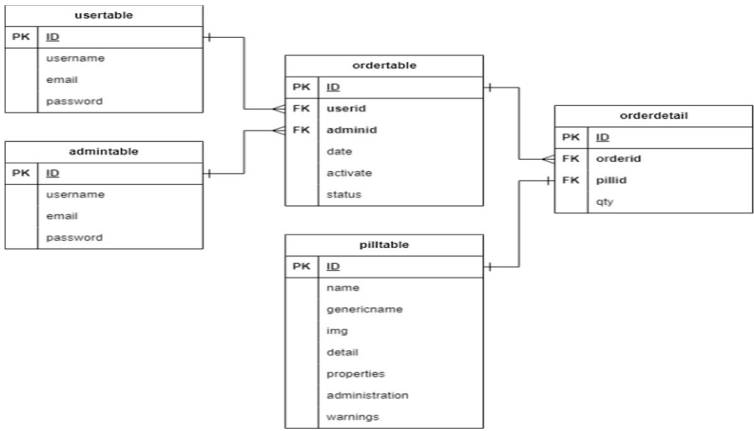


Figure 4. ER-diagram of website system.

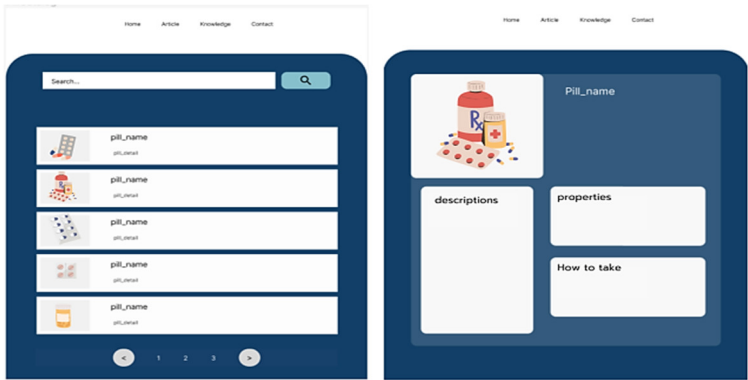


Figure 5. Displaying drug information on website.

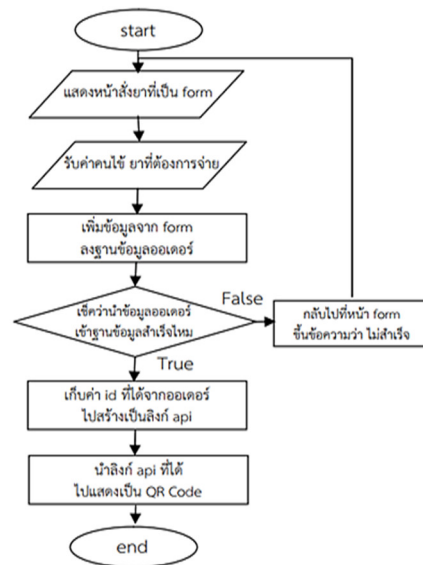


Figure 6. Flowchart of working when receiving information about the medicine used and create QR code.

2. ระบบจ่ายยาอัตโนมัติ ในระบบได้แบ่งออกมา 3 ส่วน โดยจะประกอบไปด้วยส่วนของ การออกแบบตัวของ Hardware ที่มีการนำตัวโมดูลเซอร์โว และโมดูลกลิ้งมาต่อกับ Raspberry Pi ดัง Figure 7. ที่สร้างด้วย Draw.io และในออกแบบตัวของโมเดลเครื่องจ่ายยาจะประกอบด้วยส่วนหลัก ๆ คือ กล้องที่ใช้ในการบรรจุยา แกนหมุนสำหรับบังคับส่งยา และโครงของตัวเครื่องจ่ายยา จะมีการออกแบบโมเดล 3 มิติด้วย Tinkercad

จากการพัฒนาระบบเครื่องจ่ายยาต้องการนำส่วนของ Hardware มารวมกับตัวโมเดลเพื่อที่จะสามารถใช้งานได้ ต้องมีการพัฒนาระบบในส่วนของเครื่อง โดยออกแบบกระบวนการทำงานของการจ่ายยาออกมา ดัง Figure 8. ด้วย Draw.io

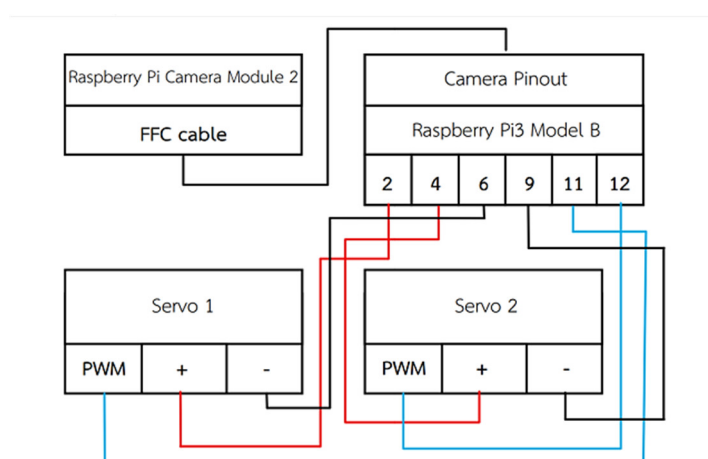


Figure 7. Diagram showing port connection of Raspberry Pi.

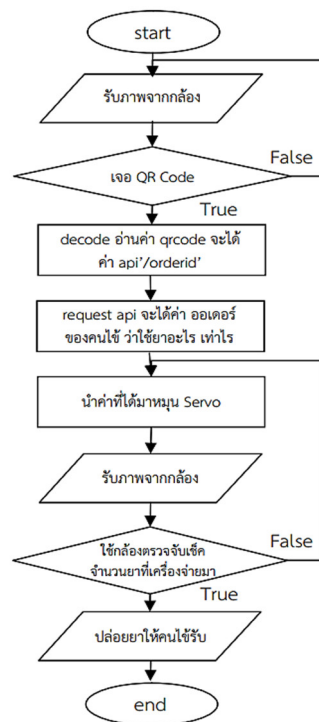


Figure 8. Flowchart of Automated Medicine Dispenser operation through QR code.

5.4 การพัฒนาระบบ

จากขั้นตอนการออกแบบ งานวิจัยนี้ได้ทำแสดงภาพรวมของระบบ และได้ทำการแบ่งระบบออกเป็น 2 ส่วน คือ 1. ระบบหน้าเว็บไซต์ และ 2.ระบบจ่ายยาอัตโนมัติ ได้มีขั้นตอนการพัฒนาระบบดังนี้

1. กำหนดปัญหาของตัวระบบ และศึกษาเกี่ยวกับข้อมูลของตัวยาตัวอย่างที่ต้องการนำมาใช้ หาความเป็นไปได้ของตัวจำลองเป็นเม็ดยา
2. วิเคราะห์ข้อมูลที่ได้จากการศึกษา เพื่อหาความต้องการของตัวระบบเว็บไซต์ พร้อมทั้งเก็บข้อมูลตัวยา รูปภาพ ขนาด รูปทรง และหาเม็ดยาจำลองตัวอย่างที่จะมาใช้ร่วมกับตัวโมเดลเครื่องจ่ายยา
3. ออกแบบฐานข้อมูลหน้าเว็บตามการวิเคราะห์ปัญหาและสิ่งที่ต้องการ โดยออกแบบกระบวนการทำงานในระบบทั้งในส่วนของการออกใบสั่งยากับการจ่ายยา ออกแบบโมเดลของเครื่องจ่ายยาให้รองรับและสามารถจ่ายเม็ดยาจำลองตามที่สั่งออกมาได้
4. พัฒนาระบบในส่วนของเว็บไซต์ และเครื่องจ่ายยาตามที่ได้ออกแบบไว้ ตามวงจรการพัฒนาระบบ (System Development Life Cycle: SDLC) โดยใช้เครื่องมือพัฒนาด้วยซอฟต์แวร์ VScode, Thonny ในการพัฒนาระบบเว็บไซต์กับระบบจ่ายยา และแพลตฟอร์ม GoogleColab ที่นำมาใช้ในการ Train โมเดล YOLOv5s ที่ใช้ตรวจจับยา

5.5 การทดสอบและปรับปรุงประสิทธิภาพระบบ

การทดสอบประสิทธิภาพของระบบเครื่องจ่ายยาเครื่องจะแบ่งเป็น 2 ส่วนคือ

1. ส่วนการทำงานของโมเดลกับ Raspberry Pi

การทดสอบประสิทธิภาพของเครื่องเนื่องจากตัวเครื่องทดสอบถูกออกแบบมาให้รองรับได้แค่ยาตัวอย่างรูปทรงเดียวเท่านั้น และสามารถรับยาได้ไม่เกิน 40 เม็ด จึงได้กำหนดการทดสอบจ่ายยาโดยสั่งให้เครื่องปล่อยยา 30 ครั้ง เพื่อที่

สามารถจะตรวจสอบได้ว่า ยาออกมาพอดีกับจำนวนที่สั่งไม่มีการจ่ายยาออกมาขาด หรือเกินที่สั่ง พร้อมบันทึกผลเป็นจำนวนยาที่เครื่องปล่อยมาได้ถูกต้อง แล้วจึงทำการทดลอง 5 ครั้ง แล้วหาเปอร์เซ็นต์ความแม่นยำจากการใช้ค่าเฉลี่ย โดยเริ่มจากการทดสอบจากโมเดล 1 เป็นโมเดลแรก ดัง Figure 9.

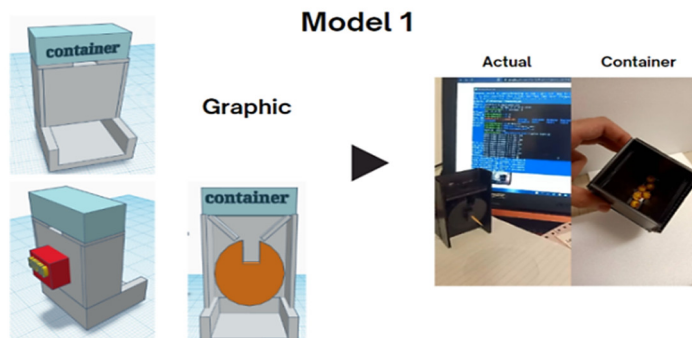


Figure 9. Design and development of Model 1.

เมื่อทดสอบประสิทธิภาพโดยวิธีข้างต้นได้อัตราความสำเร็จเพียง 38.00% เนื่องจากตัวยานั้นเบียดกันเองทำให้ยาไม่สามารถไหลลงไปสู่บริเวณแกนหมุนได้ จึงเป็นเหตุในการพัฒนาโมเดล 2 และโมเดล 2 ที่เปลี่ยนวัสดุให้มีความทนทานมากขึ้น โดยการเพิ่มขนาดวงกลมของส่วนแกนหมุนเพื่อให้ยาในภาชนะมีการเคลื่อนไหวตลอด ดัง Figure 10.

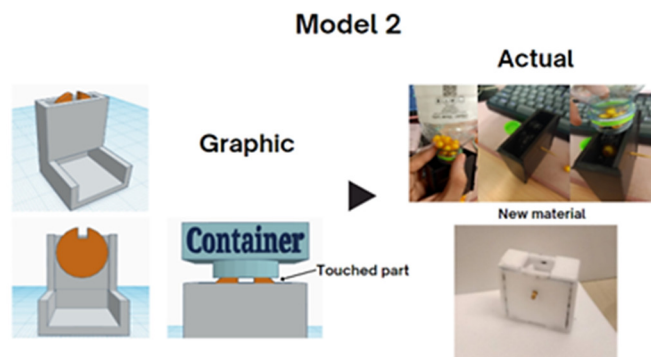


Figure 10. Design and development of Model 2.

เมื่อทดสอบประสิทธิภาพโดยใช้วิธีเดิมพบว่าโมเดล 2 และ โมเดล 2 ที่เปลี่ยนวัสดุมีความสำเร็จเพิ่มขึ้นเป็น 50.00% และ 56.00%ตามลำดับ แสดงให้เห็นถึงความต่างของประสิทธิภาพของแต่ละโมเดล แต่ว่าเนื่องจากความสำเร็จมีค่าต่ำกว่ามาตรฐาน จึงได้พัฒนา โมเดล 3 ดัง Figure 11.

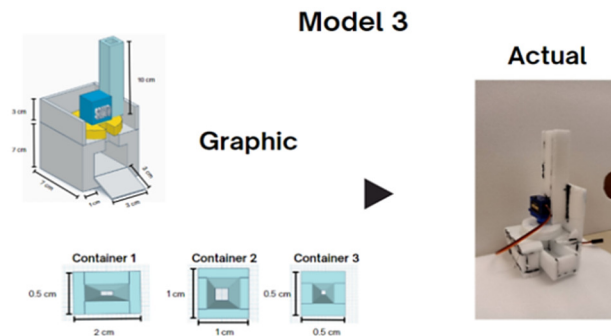


Figure 11. Design and development of Model 3.

พัฒนาโมเดลที่3 โดยเปลี่ยนรูปแบบให้แตกต่างจากเดิมมากขึ้น สามารถถอดเปลี่ยนภาชนะที่บรรจุยาได้ เพื่อรองรับขนาดต่าง ๆ โมเดล 3 มีความสำเร็จเพิ่มขึ้นเป็น 64.00%

Table 1. Test results from training the model 250 times.

Class	P	R	mAP50
ALL	0.857	0.875	0.896
Capsule	0.906	0.959	0.976
Pill1	0.759	1	0.995
Pill2	0.906	0.667	0.719

2. ส่วนการตรวจจับและนับจำนวนยาที่เครื่องจ่ายออกมา

จากตรวจจับและนับจำนวนยาด้วย โมเดล YOLOv5s ที่ผ่านการ Train ด้วยรูปภาพจำนวน 60 รูปโดยรูปภาพจะเป็นภาพถ่ายของยาทั้ง 3 กลุ่ม เป็นจำนวนกลุ่มละ 20 ภาพตามที่ได้เตรียมไว้ และการ train จำนวน 250 ครั้ง เนื่องจากผลลัพธ์การตรวจจับที่ค่อนข้างดีกว่าจำนวนรอบอื่น ๆ โดยผลลัพธ์การ Train ดัง Table 1 และได้ทำการทดสอบการตรวจสอบภาพยาด้วยโมเดล YOLOv5s ดังกล่าว กับภาพยาทั้ง 3กลุ่มจากข้อมูลรูปภาพทดสอบที่ถูกแบ่งส่วนออกมาเป็นจำนวนกลุ่มละ 10 ภาพ ได้ผลลัพธ์ดังนี้

1. ยากลุ่มที่ 1 Capsule

จากการทดสอบการตรวจจับภาพยากลุ่มที่ 1 หรือ Capsule จำนวน 10 ภาพ ผลลัพธ์ที่ได้ออกมานั้นคือ มีการตรวจพบยา Capsule ทั้งหมดที่อยู่ในภาพได้ถูกต้อง ทำให้มีความถูกต้องในการตรวจจับอยู่ 100% ตัวอย่างชุดภาพที่ใช้ในการตรวจจับดัง Figure 12.



Figure 12. Capsule detection example.

2. ยากลุ่มที่ 2 Pill1

จากการทดสอบการตรวจจับภาพยากลุ่มที่ 2 หรือ Pill1 จำนวน 10 ภาพ ผลลัพธ์ที่ได้ออกมานั้นคือ มีการตรวจพบยา Pill1 ทั้งหมดที่อยู่ในภาพได้ถูกต้อง ทำให้มีความถูกต้องในการตรวจจับอยู่ 100% ตัวอย่างชุดภาพที่ใช้ในการตรวจจับดัง Figure 13.

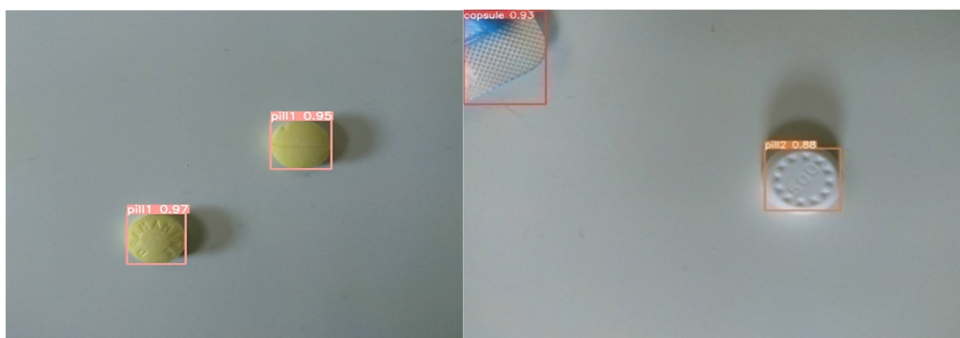


Figure 13. Pill1 detection example.

3. ยากลุ่มที่ 3 Pill2

จากการทดสอบการตรวจจับภาพยากลุ่มที่ 3 หรือ Pill2 จำนวน 10 ภาพ ผลลัพธ์ที่ได้ออกมานั้นคือ มีการตรวจพบเป็นยา Capsule 2 ภาพ และยา Pill1 จำนวน 1 ภาพ โดยส่วนใหญ่จะเป็นภาพที่มีพื้นเป็นสีขาว ทำให้มีความถูกต้องในการตรวจจับอยู่ 70% ตัวอย่างชุดภาพที่ใช้ในการตรวจจับดัง Figure 14.

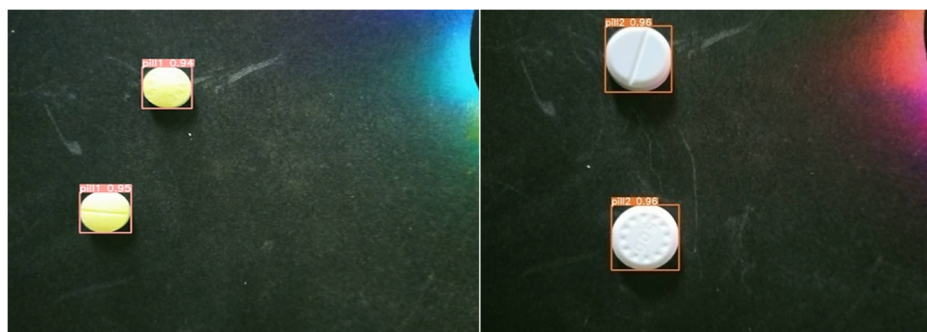


Figure 14. Pill2 detection example.

6. ผลการวิจัย (Results)

6.1 ทดสอบประสิทธิภาพการทำงานของโมเดลกับ Raspberry Pi

เมื่อพัฒนาเสร็จแล้วจึงนำไปทดสอบด้วยวิธีข้างต้นพบว่า มีอัตราความสำเร็จสูงถึง 64% ซึ่งเป็นค่าที่น่าพอใจมากขึ้น และได้บันทึกผลลงใน Table 2.

Table 2. The effectiveness of Automated Medicine Dispenser.

Test	Model 1	Model 2	Model 2 (New Material)	Model 3
1	9/30	15/30	18/30	24/30
2	12/30	18/30	18/30	21/30
3	12/30	15/30	15/30	18/30
4	15/30	12/30	15/30	18/30
5	9/30	15/300	18/30	15/30
Total (%)	38.00	50.00	56.00	64.00

จากการทดสอบประสิทธิภาพพบว่า การสื่อสารระหว่าง Raspberry Pi กับ หน้าเว็บเพจสามารถทำได้ปกติ และพบว่าโมเดลมีประสิทธิภาพมากขึ้นเมื่อเทียบกับรุ่นก่อนหน้า ปัจจัยที่ส่งผลต่อประสิทธิภาพของเครื่องจ่ายยาคือความผิดพลาดที่แบบของโมเดล และอีกหนึ่งปัจจัยที่สำคัญคือ วัสดุที่ใช้ในการพัฒนาไม่แข็งแรงเพียงพอที่จะทำงานอย่างต่อเนื่องได้ ทำให้ประสิทธิภาพลดลงเมื่อใช้งานไประยะหนึ่ง

6.2 ทดสอบประสิทธิภาพการตรวจจับและนับจำนวนยาที่เครื่องออกมา

การนำโมเดลตัว Best.pt ที่ได้มาทดสอบการตรวจจับกับชุดรูปภาพยาสำหรับทดสอบจำนวน 30 รูป โดยแบ่งออกเป็น 3 กลุ่ม กลุ่มละ 10 รูป มีผลตาม Table 3.

Table 3. Drug image detection test results

Test results	Capsule	Pill1	Pill2
Correct detection	10	10	7
Incorrect detection	0	0	2
No detection	0	0	0
Accuracy	100%	100%	70%
Number of samples	10	10	10
Total accuracy	90%		

7. สรุปผลการวิจัย (Conclusion)

ผลการวิจัยเรื่อง เครื่องจ่ายยาอัตโนมัติและตรวจสอบความถูกต้องด้วยการใช้ปัญญาประดิษฐ์ สามารถสรุปผลการวิจัยได้ ดังนี้

1. ผลการพัฒนากระบวนการเครื่องจ่ายยาอัตโนมัติ พบว่า การสื่อสารระหว่าง Raspberry Pi กับหน้าเว็บไซต์สามารถทำงานได้ปกติ และมีประสิทธิภาพของโมเดล 1, โมเดล 2, โมเดล 2 (วัสดุใหม่) และโมเดล 3 อยู่ที่ 38.00%, 50.00%,



56.00% และ 64.00% ตามลำดับ ดังนั้น จึงสรุปได้ว่า ระบบเครื่องจ่ายยาอัตโนมัติ สามารถใช้งานได้อย่างมีประสิทธิภาพ มีความคลาดเคลื่อนน้อยและลดเวลาที่ใช้ในการเตรียมยาน้อยลง

2. ผลการพัฒนาระบบตรวจสอบความถูกต้องด้วยการใช้ปัญญาประดิษฐ์ พบว่า ระบบตรวจสอบความถูกต้องสามารถทำงานได้อย่างมีความแม่นยำ และมีผลรวมความแม่นยำของโมเดลสูงถึง 90% จากการทดลองตรวจจ่ายยาทั้ง 3 ชนิด คือ Capsule, Pill1 และ Pill2 อย่างละ 10 เม็ด รวม 30 เม็ด พบว่ามีความแม่นยำของแต่ละกลุ่มอยู่ที่ 100%, 100% และ 70% ตามลำดับ ดังนั้น จึงสรุปได้ว่าระบบตรวจสอบความถูกต้องสามารถทำงานได้ดี ส่งผลให้อัตราการเกิดความคลาดเคลื่อนทางยาลดลงเพราะมีการตรวจสอบอีกครั้งดัง Table 4.

Table 4. Confusion Matrix of three types of drug.

Actual	Predict			
	9/30	Capsule	Pill1	Pill2
Capsule		10	0	2
Pill1		0	10	1
Pill2		0	0	7

8. อภิปรายผลการวิจัย (Discussion)

จากการศึกษา เครื่องจ่ายยาอัตโนมัติและตรวจสอบความถูกต้องด้วยการใช้ปัญญาประดิษฐ์ มีประเด็นการอภิปรายผลดังนี้

1. ผลการพัฒนาระบบเครื่องจ่ายยาอัตโนมัติ พบว่า โมเดลจะมีประสิทธิภาพมากขึ้นเมื่อเทียบกับรุ่นก่อนหน้า และพบว่าปัจจัยที่ส่งผลต่อประสิทธิภาพของเครื่องจ่ายยาคือความผิดพลาดที่การออกแบบโมเดลและวัสดุที่ใช้ในการพัฒนา ไม่แข็งแรงเพียงพอที่จะทำงานอย่างต่อเนื่อง ทำให้ประสิทธิภาพลดลงเมื่อใช้งานไประยะหนึ่งได้ผลการพัฒนาระบบตรวจสอบความถูกต้องด้วยการใช้ปัญญาประดิษฐ์ พบว่าระบบสามารถตรวจสอบกลุ่มยาแต่ละชนิดได้ปกติยกเว้น Pill2 เนื่องจากในการบวนการตรวจสอบของปัญญาประดิษฐ์จะมีการเร่งค่าความสว่าง (Brightness) ของรูปภาพก่อนที่จะทำการตรวจสอบทำให้ Pill1 และ Pill2 มีความใกล้เคียงกัน

2. ผลการใช้เครื่องจ่ายยาเพื่อลดความคลาดเคลื่อนทางยา พบว่า เครื่องจ่ายยานี้มีความคลาดเคลื่อนทางยาอยู่ที่ 23% เนื่องจากโมเดลของตัวเครื่องจ่ายยา แต่ว่าเมื่อรวมกับระบบตรวจสอบความถูกต้อง ทำให้ลดอัตราการเกิดความคลาดเคลื่อนอยู่ที่ 11% ซึ่งน้อยกว่าความคลาดเคลื่อนทางยาที่เกิดจากมนุษย์ตามงานวิจัยของ Mumford et al. (2024) ที่บอกว่าความคลาดเคลื่อนทางยาอยู่ที่ 74%

9. ข้อเสนอแนะงานวิจัย (Recommendation)

จากผลการทดลองข้างต้นทำให้ทำให้ได้ข้อสรุปว่าตัวผลิตภัณฑ์ชิ้นนี้สามารถใช้งานระบบการออกใบสั่งยา ผ่านระบบหน้าเว็บไซต์ ได้ถูกต้องแต่ผลที่ได้ยังไม่น่าพอใจ สาเหตุหลักต่อความถูกต้องในการทำงานนั้นมาจาก การออกแบบที่ไม่ดีพอ และวัสดุที่เลือกใช้ไม่เหมาะสม หากต้องการที่จะเพิ่มประสิทธิภาพในการทำงานนี้ คือ ต้องออกแบบผลิตภัณฑ์ใหม่และใช้วัสดุชนิดอื่น ๆ ในการสร้าง เพื่อให้การทำงานเป็นไปอย่างราบรื่นและมีความทนทานมากขึ้นเนื่องด้วยงานวิจัยชิ้นนี้เป็นการพัฒนาระบบเครื่องจ่ายยาอัตโนมัติ สามารถนำผลิตภัณฑ์ชิ้นนี้ใช้ในมางานด้านการจ่ายยาในโรงพยาบาลได้ เช่น การนำเครื่องจ่ายยาอัตโนมัติไปใช้เป็นจ่ายยาสามัญทั่วไปให้กับคนไข้ที่มาเข้ารับบริการในโรงพยาบาล เพื่อเป็นการลดเวลาในการต่อคิวรับยา และช่วยเพิ่มความรวดเร็วในการบริการจ่ายยาของโรงพยาบาล

10. เอกสารอ้างอิง (References)

- ABB. (2020, April 20). *What is Deep Learning?*. <https://new.abb.com/news/detail/58004/deep-learning>. (In Thai)
- Al-Fuqaha, A., Guizani, M., Mohammadi, M., Aledhari, M., & Ayyash, M. (2015). Internet of Things: A Survey on Enabling Technologies, Protocols, and Applications. *IEEE Communications Surveys & Tutorials*, 17(4), 2347–2376. <https://doi.org/10.1109/COMST.2015.2444095>.
- Almotiri, S. H., Khan, M. A., & Alghamdi, M. A. (2016). Mobile Health (m-Health) System in the Context of IoT. *2016 IEEE 4th International Conference on Future Internet of Things and Cloud Workshops (FiCloudW)*, 39–42. <https://doi.org/10.1109/W-FiCloud.2016.24>.
- Amazon Web Services. (n.d.). *What is API (Application Programming Interfaces)?*. <https://aws.amazon.com/th/what-is/api>. (In Thai)
- Anderson, B., & Nicholson, B. (2022, June 12). SQL vs. NoSQL Databases: What's the Difference? *IBM Blog*. <https://www.ibm.com/blog/sql-vs-nosql>.
- Anupama, A. H., Srilekha, G., & Uma Priya, G. (2020). Review on Artificial Intelligence (AI) in Drug Dispensing and Drug Accountability. *International Journal of Pharmaceutical Sciences Review and Research*, 61(2), 32-35. <https://globalresearchonline.net/journalcontents/v61-2/07.pdf>.
- Barker, K. N., Pearson, R. E., Hepler, C. D., Smith, W. E., & Pappas, C. A. (1984). Effect of an Automated Bedside Dispensing Machine on Medication Errors. *American Journal of Health-System Pharmacy*, 41(7), 1352–1358. <https://doi.org/10.1093/ajhp/41.7.1352>.
- Bu, F., Sun, H., Li, L., Tang, F., Zhang, X., Yan, J., Ye, Z., & Huang, T. (2022). Artificial Intelligence-based Internet Hospital Pharmacy Services in China: Perspective Based on a Case Study. *Frontiers in Pharmacology*, 13, 1027808. <https://doi.org/10.3389/fphar.2022.1027808>.
- Chen, Z., Wu, R., Lin, Y., Li, C., Chen, S., Yuan, Z., Chen, S., & Zou, X. (2022). Plant Disease Recognition Model Based on Improved YOLOv5. *Agronomy*, 12(2), 365. <https://doi.org/10.3390/agronomy12020365>.
- Dayananda, P., & Upadhyay, A. G. (2024). Development of Smart Pill Expert System Based on IoT. *Journal of The Institution of Engineers (India): Series B*, 105(3), 457–467. <https://doi.org/10.1007/s40031-023-00956-2>.
- de Bienassis, K., Esmail, L., Lopert, R., & Klazinga, N. (2022). *The Economics of Medication Safety: Improving Medication Safety through Collective, Real-time Learning*. OECD Publishing. <https://doi.org/10.1787/9a933261-en>.
- Gargioni, L., Fogli, D., & Baroni, P. (2024). A Systematic Review on Pill and Medication Dispensers from a Human-Centered Perspective. *Journal of Healthcare Informatics Research*, 8(2), 244–285. <https://doi.org/10.1007/s41666-024-00161-w>.
- Kim, J., Kwon, H., Kim, J., Park, J., Choi, S.-U., & Kim, S. (2022). PillGood: Automated and Interactive Pill Dispenser Using Facial Recognition for Safe and Personalized Medication. *Proceedings of the*

Thirty-First International Joint Conference on Artificial Intelligence, 5920–5923.

<https://doi.org/10.24963/ijcai.2022/854>.

- Mumford, V., Raban, M. Z., Li, L., Fitzpatrick, E., Woods, A., Merchant, A., Badgery-Parker, T., Gates, P., Baysari, M., Day, R. O., Ambler, G., Dalla-Pozza, L., Gazarian, M., Gardo, A., Barclay, P., White, L., & Westbrook, J. I. (2024). Developing a Process to Measure Actual Harm from Medication Errors in Paediatric Inpatients: From Design to Implementation. *British Journal of Clinical Pharmacology*, 90(7), 1615–1626. <https://doi.org/10.1111/bcp.16052>.
- Nasir, Z., Asif, A., Nawaz, M., & Ali, M. (2023). Design of a Smart Medical Box for Automatic Pill Dispensing and Health Monitoring †. *INTERACT 2023*, 7. <https://doi.org/10.3390/engproc2023032007>.
- Niu, Y., Wang, L., Yu, Z., Huang, J., Huang, B., & Su, Y. (2023). Vision-based Automatic order Check Method for Online Medicine Dispensing Cabinet under Incomplete Data. *Engineering Applications of Artificial Intelligence*, 123, 106204. <https://doi.org/10.1016/j.engappai.2023.106204>.
- Ravi, D., Wong, C., Deligianni, F., Berthelot, M., Andreu-Perez, J., Lo, B., & Yang, G.-Z. (2017). Deep Learning for Health Informatics. *IEEE Journal of Biomedical and Health Informatics*, 21(1), 4-21. <https://doi.org/10.1109/JBHI.2016.2636665>.
- Shen, D., Wu, G., & Suk, H.-I. (2017). Deep Learning in Medical Image Analysis. *Annual Review of Biomedical Engineering*, 19(1), 221–248. <https://doi.org/10.1146/annurev-bioeng-071516-044442>.
- Watson, D. (2018, July 24). *Introduction to Raspberry Pi 3 B+*. *The Engineering Projects*. <https://www.theengineeringprojects.com/2018/07/introduction-to-raspberry-pi-3-b-plus.html>.
- Xu, X., Zhang, X., & Zhang, T. (2022). Lite-YOLOv5: A Lightweight Deep Learning Detector for On-Board Ship Detection in Large-Scene Sentinel-1 SAR Images. *Remote Sensing*, 14(4), 1018. <https://doi.org/10.3390/rs14041018>.
- Zanella, A., Bui, N., Castellani, A., Vangelista, L., & Zorzi, M. (2014). Internet of Things for Smart Cities. *IEEE Internet of Things Journal*, 1(1), 22–32. <https://doi.org/10.1109/JIOT.2014.2306328>.
- Zheng, Y., Rowell, B., Chen, Q., Kim, J. Y., Kontar, R. A., Yang, X. J., & Lester, C. A. (2023). Designing Human-Centered AI to Prevent Medication Dispensing Errors: Focus Group Study With Pharmacists. *JMIR Formative Research*, 7, e51921. <https://doi.org/10.2196/51921>.

Comparing a Thai Words Segmentation Methods in the LST20 Dataset

Krittapol Damrongkamoltip¹, Khatcha Ruenlek², Wasit Limprasert^{3,*}, and Prachya Boonkwan^{4,*}

Received: 26 March 2024;

Revised: 13 May 2024;

Accepted: 17 May 2024;

Published: 12 August 2024;

Abstract

In this era of globalization where information is widely available, organizations are increasingly placing importance on using information to enhance their business. Although data is easily available, there are still challenges in natural language processing tasks, especially, the division of Thai words that lacks clarity of word boundaries, etc. This makes it difficult to identify the word groups in a sentence appropriately. Therefore, this study focuses on evaluating the performance of the word segmentation method including the Dictionary use and learning from data using evaluation of word segmentation in six techniques are important goals for the verification of the literal level accuracy and processing time of each method and technique, by the LST20 dataset contains 3,745 documents and covers 15 news categories in results show a more efficient way to learn from data.

Keywords: Thai Words, Segmentation Methods, LST20 Dataset.

1. Introduction

In the age of social media explosion, data has become vital for organizations seeking to gain a competitive advantage and optimize management processes. Text data, prevalent across social platforms, fuels a process called Natural Language Processing (NLP) that unlocks valuable insights. NLP plays a critical role in a wide range of real-world applications, especially those involving user interaction. One of the major challenges in language processing is word segmentation, also known as word tokenization. This process involves breaking down written language into distinct units. While some languages, like English, Spanish, and French, benefit from clear word boundaries and relatively fixed structures, others like Chinese, Thai, and Vietnamese present complexities due to the absence of explicit word boundaries. In

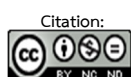
¹ Student, Data Science and Innovation, College of Interdisciplinary Studies, Thammasat University, Pathum Thani 12121, Thailand; Email: krittapol.dam@dome.tu.ac.th

² Student, Data Science and Innovation, College of Interdisciplinary Studies, Thammasat University, Pathum Thani 12121, Thailand; Email: khatcha.rue@dome.tu.ac.th

^{3,*} Assistant Professor, Dr., Data Science and Innovation, College of Interdisciplinary Studies, Thammasat University, Pathum Thani 12121, Thailand; Email: wasit@tu.ac.th

^{4,*} Lecturer, Dr., Language and Semantic Technology Research Team, National Electronics and Computer Technology Center, Pathum Thani 12121, Thailand; Email: prachya.boonkwan@nectec.or.th

*Corresponding authors: Wasit Limprasert (wasit@tu.ac.th); Prachya Boonkwan (prachya.boonkwan@nectec.or.th)



these languages, words can span multiple syllables, and context often dictates where a word begins and ends. Thai words, for example, often lack clear delimiters and can encompass multiple syllables.

Accurate segmentation is paramount for NLP applications as it directly impacts tasks such as part-of-speech tagging, named entity recognition, and machine translation. Improved segmentation not only enhances the accuracy of these applications but also facilitates deeper linguistic analysis and more effective information retrieval. Therefore, understanding the complexities of Thai word segmentation and achieving precise segmentation are essential for advancing NLP research and applications in the Thai language. Various tools, ranging from dictionary-based methods to learning-based segmentation and sub word tokenization, offer solutions to address these challenges.

This study focuses on Thai language processing, aiming to evaluate the accuracy and efficiency of different segmentation methods. Our contribution lies in establishing a benchmark to determine the most suitable methods for diverse data contexts. We aim to provide valuable insights for selecting the most appropriate segmentation methods, ultimately reducing resource demands, and ensuring proper functionality across various data contexts. To achieve this, we compare six prevalent Thai word segmentation methods – Newmm, Colloc, Longest Matching, Attacut, Deepcut, and XLM-RoBERTa. We leverage the LST20 corpus (Boonkwan, 2020; National Electronics and Computer Technology Center, 2022) curated by Thailand's NECTEC, which comprises 3,745 annotated documents across 15 news genres. This extensive dataset provides a rich resource for in-depth exploration of Thai word segmentation methods.

2. Literature Review

Word segmentation in the Thai language has many challenges. For example, the lack of space delimiters, unlike other languages such as English, which has spaces between words, and the use of compound words, which are combinations of multiple words or syllables. (Haruechaiyasak et al., 2008) For example, the classic case "ตากลม" can be segmented in two different ways: "ตาก-ลม" (air dry) and "ตา-กลม" (round eyes) or "แม่น้ำ" (river), which is composed of "แม่" (mother) and "น้ำ" (water). Previous studies have explored different methodologies to tackle this issue, offering valuable insights into the evolution of approaches over time.

The existing literature on Thai word segmentation has provided valuable insights into different methodologies. For instance, a study titled "A Comparative Study on Thai Word Segmentation Approaches" conducted in 2008 compared dictionary-based (DCB) approaches to machine learning-based (MLB) approaches using n-gram transformations from the ORCHID corpus (Noyunsan et al., 2014). However, this study primarily focused on performance metrics, neglecting the crucial aspect of processing time, which is often a constraint in real-world applications with deadlines. Expanding on this, "A Multi-Aspect Comparison and Evaluation on Thai Word Segmentation Programs" in 2014 evaluated six Thai word segmentation programs: Libthai, Swath, Wordcut, CRF++, Thaisemantics, and Tlexs, using the BEST corpus 2014 developed by NECTEC (Aroonmanakun, 2002). This research not only assessed accuracy but also

considered processing time. Each program varied in its implementation environment, with differences in programming languages and support systems. While these studies contributed significantly to understanding Thai word segmentation, they also revealed certain limitations. For instance, each segmentation program required specific installations and environments, which could be cumbersome for users. Additionally, the rapid evolution of programming languages and frameworks has shifted developers' preferences towards specialization in a few languages rather than mastering multiple ones. To address these challenges and advance the field, our proposed research aims to compare various Thai word segmentation methods within a unified environment. By leveraging Python and utilizing the latest Thai language corpus, LST20, we intend to conduct a comprehensive comparison focusing on both accuracy and preprocessing time. This approach ensures consistency in evaluation metrics and simplifies the implementation process for users, aligning with contemporary trends in programming specialization and efficiency.

Regarding methodologies, two common approaches have been widely employed: the dictionary-based approach and the learning-based approach. The dictionary-based approach uses a dictionary to parse and segment text into words. This approach has a few popular selection methods, such as the longest-matching, which attempts to match a set of characters into the longest possible word (Poonwarawan, 1986). Another popular method is the maximum-matching, which finds every possible way to segment text and then selects the way with the fewest words (Virach, 1993). However, in this paper, we don't directly use the maximum-matching. Instead, we use the engine "newmm" from the library "PyThaiNLP." This engine makes use of the maximum-matching method and the Thai Character Cluster (TCC). The Thai Character Cluster uses the concept of a character cluster, which is a unit smaller than a syllable but larger than a character and matches them with defined rules (Theeramunkong et al., 2000). Additionally, we test the maximum collocation approach, which uses the idea of collocation strength between syllables to form words. Collocation can also be referred to as the co-occurrence of syllables, i.e., two syllables that are part of a word will have higher collocation strength. They use a dictionary to match all possible words from the syllables and then use collocation strength to select the best segmentation with the highest collocation strength.

However, advancements in machine learning have led to the learning-based approach, the word segmentation problem can be viewed as a binary-classification problem, trying to classify if a character is a word-beginning character (B) or a word-inning character (I). We compare three models for this problem: DeepCut (Kittinaradorn et al., 2019), AttaCut (Chormai et al., 2019), and an XLM-RoBERTa-based POS tagging model (De Vries et al., 2022). For the DeepCut model, they use a Deep Neural Network to perform the binary classification task. The Convolutional Neural network was trained with NECTEC's BEST corpus (Kittinaradorn et al., 2019). In the case of the AttaCut model, it's constructed using CNNs like the DeepCut model. However, there have been adjustments made to the convolutional layers, specifically focusing on minimizing overlap between them. As a result of this modification, the execution time is reduced compared to the DeepCut model. Additionally, AttaCut incorporates syllable knowledge through a

process known as syllable embedding, where characters within the same syllable share identical syllable embeddings (Chormai et al., 2019).

Lastly, the XLM-RoBERTa-based POS tagging model. The model we use is called XLM-RoBERTa base Universal Dependencies v2.8 POS tagging (De Vries et al., 2022). This model is fine-tuned from the pre-trained XLM-RoBERTa model (Conneau et al., 2020) for part-of-speech tagging using many different combinations of training and testing languages. The specific one we use was trained on the Hebrew language. As we mentioned, this model was originally trained for POS tagging, but it can also be used for word segmentation as it also has word boundaries as output.

3. Research Methodology

The experimental process consists of three phases: data preparation, data processing, and data measurement. The process commences with downloading the dataset through the openD portal (National Electronics and Computer Technology Center, 2022). Subsequently, data preprocessing is conducted to transform it into a usable format such as CSV using the DataFrame from the Pandas package. Concurrently, an environment for the experiment is set up, involving the download and installation of necessary tools, including Pip for acquiring essential packages and Python for running a script. After establishing the environment and completing the data preprocessing step, the second process consists of two parts, focusing on data measurement dimensions: accuracy and time processing. The transformed data is processed in a function implemented to work with six word segmentation methods: Newmm, Deepcut, Attacut, Longest, TLTK, and XML-Roberta. Almost all these methods are downloaded via Pip, except XML-Roberta, which is processed using an API pipeline implemented from Hugging Face. The results from each function are saved and exported in CSV format for use in subsequent stages. In the time processing stage, 300 records of mock-up data (De Vries et al., 2022) are used to examine the time processing of each word segmentation method. The mock-up data is divided into three equally sized datasets. Analogous to the accuracy processing, the data is passed through the word segmentation function. In addition to parsing the data before parsing the input data, the time is saved using the time module. Also, after parsing, the time is recorded as well. Finally, in the data measurement process, the results from the second process are compared with the ground truth answer. Accuracy is calculated using three metrics: Precision, Recall, and F1 score, employing a confusion matrix approach. Regarding the processing time is determined by finding the difference between the start and end parsing times.

5.1 Data Preparation

Our experiment began with the download of the LST20 corpus dataset from the openD portal (National Electronics and Computer Technology Center, 2022). Following the download, we set up the processing environment, which included the installation of Python and Pip. Five essential packages Pythainlp, Deepcut, Attacut, TLTK, and Request which needed to be downloaded to correspond with the selected word segmentations. We ensured that we installed the latest versions of all tools available at that time, namely Pythainlp v4.0.2, Deepcut v0.7.0, AttaCut v1.0.6, and TLTK v1.8.

Once the environment was configured, we imported the LST20 corpus from text files into a pandas dataframe which each file consists of many words that make up a sentence. In one row, there will be five components: word boundaries, POS tagging, named entities, clause boundaries, and sentence boundaries. If a word is a space, they use "_" as a symbol for space. This was achieved by utilizing the built-in function open() to open files from their path, retrieving content through the read() method, and then saving the content to dataserie and dataframe, respectively. The dataframe comprised two columns: the original sentences, concatenated from tokens in text files with join function, And the ground truth is representing the original data within the text files. To avoid processing errors, we consistently saved all outputs during the process as CSV files from the pandas dataframe, using the to_csv() function and setting the "encoding" parameter to "utf-8 - sig" to ensure readability when opened by other applications.

5.2 Data Processing

The processing began by importing the necessary packages mentioned above into the Python editor. We proceeded to tokenize the first column of original sentences until all sentences were processed, recording the start and end times for each sentence. This process is iterated in a loop for all selected word segmentations. During the parsing of sentences, the process unfolded as follows: first, we imported the necessary packages into the editor. Next, default values were set for each parameter of each tokenizer. For Attacut, the default model is "attacut-sc," which represents the best-performing model trained by the Attacut package. Regarding Deepcut, we utilize it directly from the Deepcut package by importing the tokenizer as follows: "from deepcut import tokenize as deepcut_tokenizerparameters," without specifying any additional parameters. Similarly, for TLTK, we employ it directly from the TLTK package using the following import statement: "from tltk.nlp import word_segment as tltk_tokenizer," where the word segmentation method utilized is based on the maximum collocation approach, denoted as 'colloc' by default. For both Newmm and Longest matching, we utilize them via the Pythainlp class named "from pythainlp.tokenize import word_tokenize as pythai_default_tokenizer" This class encompasses four parameters: engine, custom_dict, keep_whitespace, and join_broken_num. These parameters default to certain values, with custom_dict based on the Thai National Corpus (TNC) and keep_whitespace and join_broken_num set to True. The only parameter that differs among these methods is the engine parameter, which specifies the model or tokenizer object to be used. It is set to "newmm" for Newmm and "longest" for the longest matching approach. Before the iteration commenced, we need to create the list to store the output from the loop processed, assumed named as (y_tokenized_lst). The iteration commenced by taking a sentence from the first column of the imported dataframe or the full sentences from each file of LST20, stored in a temporary variable assumed named (X_text). Then, this temporary variable (X_text) was passed into the tokenizer, and the resulting tokenized output was stored in y_tokenized_lst. This process was repeated until all sentences from the first column were tokenized, and the entire process was replicated for other word segmentations.



Citation:

Damrongkamoltip, K., Ruenlek, K., Limprasert, W., & Boonkwan, P. (2024). Comparing a Thai Words Segmentation Methods in the LST20 Dataset. *Journal of Computer and Creative Technology*, 2(2), 61-70. <https://doi.org/10.14456/jcct.2024.7>.

And the last method XLM-Roberta model of sub word tokenization approach required different steps. First, we signed up for a Hugging Face account to create an API key token. Then, we imported the Request package into the editor and set variable values, including API_URL (the URL of the interested Hugging Face model API) and Headers ("Authorization": f"Bearer {'api_key'}"). Following this, we created a function for sending requests to the server using the POST method and set the POST parameters to "requests.post(API_URL, headers=headers, json=payload)" with payload representing the input data from the function. The output from the return method was set to JSON format for flexibility and ease of application. Subsequently, we executed the same process, passing the variable (x) storing sentences into the XLM-Roberta model function for tokenization. The resulting tokenized output was stored in another variable. Since the output from the XLM-Roberta model was in JSON format, we extracted only the required tokenized words from the JSON output using a loop and accessing the data value by the key named 'word.' These words were stored in a list, and the process was repeated until all sentences were tokenized.

Once the word segmentation was complete, we transformed the results from the tokenized process into a usable data type: a list of tokens. TLTK was the only word segmentation that returned results as strings but separated each word inside the sentence using "|". If it couldn't tokenize words, it returned the output as "<Fail>...</Fail>" or "\xa0," constraints that needed to be eliminated. After transforming the outputs, the next step was to calculate accuracy at the character-based level using the confusion matrix approach.

The confusion matrix serves as a table to evaluate the performance of a classification algorithm, employing four key elements: True Positive, False Positive, True Negative, and False Negative. These elements are variables in the formula for calculating metrics like precision, recall, and F1 score. Throughout this process, we conducted a character-based comparison between a list of tokenized tokens and the ground truth. Each character from both lists was scrutinized against one another, utilizing the last character of the token and its position to index the confusion matrix elements. To calculate accuracy, we transformed the last character from all tokens in the lists of all segments into new special characters absent in any of the sentences. Simultaneously, the last character of every token in the ground truth lists underwent a similar conversion. Following this conversion, all tokenized token elements in the list were concatenated into a sentence. Characters in the list of tokenized tokens were then classified based on their position in confusion matrix elements. The classification process encompassed storing the tokenized sentence and the ground truth sentence in variables, ensuring both sentences were of the same length. These sentences were then taken to a loop function to iterate through each character. Rules for classification were established: if a ground truth character matched a previously converted special character in the tokenized sentence, the True Positive value increased accordingly; if the ground truth character was a special character not matched in the tokenized sentence, the False Negative value increased; if the ground truth character was a different character and matched the position in the tokenized sentence, the True Negative value increased; for all other cases, the False Positive value

increased. This comprehensive process was repeated for every sentence from every word segmentation method, comparing them with the ground truth.

5.3 Time Processing

To measure the time processing, we randomly selected 300 sentences from a data frame containing LST20 data and divided them into three sets, each comprising 100 sentences. The execution process followed these steps: First, we imported the time package into the editor. Then, for each set of data, we brought it into the tokenizer and recorded the start time before tokenization using the `time.time()` function. Subsequently, we tokenized the sentences and recorded the end time after the words were cut, storing the time before and after parsing words separately in a list. We repeated this process until all sentences were tokenized and then repeated it again with other data sample sets. This entire procedure was continued until every word segmentation method had been thoroughly tested.

5.4 Data Measurement

Following the process of tokenization and subsequent tabulation of the elements constituting the confusion matrix, two pivotal dimensions were meticulously gauged: accuracy and time. In delving into accuracy, a trifecta of critical metrics—precision, recall (also known as sensitivity rate), and F1-score—stood as the pillars of assessment. Each metric was calculated for every sentence, and the results were incorporated into the dataframe. The average was then calculated for each word segmentation using the mean function, providing the arithmetic average through the column. Regarding accuracy measurement, three metrics were used. Precision measures the accuracy of positive predictions made by the model, calculating the ratio of true positive predictions to the total number of positive predictions made by the model (both correct and incorrect). Recall measures the ability of the model to correctly identify all relevant instances, calculating the ratio of true positive predictions to the total number of actual positive instances in the dataset. F1 Score represents the harmonic mean of precision and recall, providing a balanced measure between the two metrics. When it came to measuring time, the process involved identifying the gap between when a sentence started and when it ended. This process aimed to determine the average time for each word breakdown, employing the same method used for accuracy assessment. By utilizing a mathematical mean function, which calculates an average, the total time for all word divisions was aggregated and then divided by the number of divisions, resulting in a number representing the average time taken.

4. Results and Conclusion

The results of link code and dataset were updated on the The Github website (<https://github.com>) (Ruenlek & Damrongkamoltip, 2023) and the results of comparing were as follows Table 1. and Table 2.

Table 1. A table showed the accuracy of six models across three metrics: precision, recall, and F1-score.

Measurement	Newmm	Colloc	Longest	Deepcut	Atta cut	xlm-roberta
Precision	0.94	0.94	0.12	0.94	0.93	0.13
Recall	0.85	0.89	0.31	0.81	0.90	0.27
F1 Score	0.89	0.91	0.17	0.87	0.92	0.17

Table 2. A table shows the processing time of six models across three sample dataset each 100 records.

Segmentation Methods	Processing Times			Average Processing Time
	1st Test	2nd Test	3rd Test	
Newmm	00.05	00.08	00.04	00.05 s
Colloc	101.47	30.13	21.90	51.17 s
Longest	12.74	330.98	06.61	116.78 s
Deepcut	04.47	02.70	01.68	02.95 s
Attacut	00.30	00.22	00.15	00.22 s

The results from the tables above demonstrate that both dictionary-based approach and learning-based approach have high accuracy methods and low accuracy methods. However, when considering processing time, the results above indicate that the learning-based approach processes faster than the dictionary-based approach.

4.1 Accuracy Dimension

Best Models (High Accuracy): Attacut and Colloc exhibit high accuracy across all three metrics—precision, recall, and F1-score. These models consistently achieve high values in all accuracy metrics, making them suitable choices for accurate word segmentation.

1) Precision: Deepcut, Attacut, and Newmm are the best, all scoring over 0.9.

2) Recall: Attacut stands out with 0.90, followed by Colloc (0.89) and xlm-roberta (0.27) is lowest recall score.

3) F1-Score: Attacut, and Colloc perform the best, each achieving over 0.90.

Average Models (Moderate Accuracy): Newmm and Deepcut perform reasonably well, achieving moderate accuracy levels across all metrics. While its performance is not as high as Attacut and Colloc, it still offers a decent balance between precision, recall, and F1-score.

Worst Models (Low Accuracy): Longest and xlm-roberta perform poorly in terms of accuracy. They have significantly lower precision, recall, and F1-score values compared to other models, indicating their inability to accurately segment words in the given context.

4.2 Processing Time Dimension

Fastest Models (Low Processing Time): Attacut, Newmm, and xlm-roberta are the fastest models in terms of processing time, with Newmm being the quickest overall. These models require very little time to process each sentence, making them suitable for applications where speed is a critical factor.

Moderate Processing Time: Deepcut has moderate processing times, with Deepcut being slightly faster. They strike a balance between accuracy and processing speed, making them viable options for applications where accuracy is essential, but speed is also a concern.

Slowest Model (High Processing Time): Longest and Colloc have a significantly longer processing time compared to other models, making it the slowest option. Its accuracy is also low, making it less preferable for most practical applications.

Best Model Selection: If the primary concern is accuracy and precision in word segmentation tasks, Attacut, Deepcut, and Newmm are the top choices respectively. They consistently outperform other models in accuracy metrics, making them suitable for applications where precise word boundaries are crucial. **Consideration for Specific Use Cases:** If the task requires very fast processing with a compromise on accuracy, Attacut and Newmm are the quickest options. However, it's essential to assess the specific use case and determine whether the sacrificed accuracy aligns with the application's requirements. **Avoiding the Slowest and Least Accurate Model:** Longest is both the slowest and least accurate option. Unless there is a specific use case where its characteristics align with the requirements, it is advisable to avoid using Longest due to its low accuracy and high processing time.

5. Discussion

The study explored the performance evaluation of various Thai word segmentation methods, shedding light on their accuracy and processing time. The findings underscore the importance of tailoring segmentation methods to specific use cases. For tasks prioritizing high accuracy, options like Attacut and Colloc emerge as frontrunners. Notably, learning-based approaches like Attacut or Deepcut demonstrate prowess in handling data containing new or informal words. Despite being trained mainly on formal data, their neural network-based architecture enables them to delineate word boundaries more effectively compared to segmentation methods based on dictionary approaches. Conversely, for swift processing, Attacut and Newmm stand out, albeit with a compromise on accuracy. However, when dealing with tasks involving formal data or requiring sentence segmentation into formal words, the dictionary-based approach emerges as a strong candidate. In this method, all words to be segmented are stored in a dictionary. Conversely, if the task involves detecting or identifying slang words or technical terms, this approach should be at the top of the list of potential solutions. This nuanced understanding enables practitioners to make informed decisions aligning with their application requirements.

Furthermore, the study prompts consideration of future avenues. Testing with diverse datasets, especially from platforms like TikTok and Twitter, offers a glimpse into real-world language complexities, encompassing slang, informal language, and varied topics. Exploring hybrid segmentation methods or integrating transformer models could potentially enhance accuracy and adaptability to evolving linguistic landscapes.

6. References

- Aroonmanakun, W. (2002). Collocation and Thai Word Segmentation. *Proceedings of SNLP-Oriental COCOSA*, 68-75. https://www.academia.edu/21349075/Collocation_and_Thai_Word_Segmentation.
- Boonkwan, P., Luantangsrisk, V., Phaholophinyo, S., Kriengkiet, K., Leenoi, D., Phrombut, C., Boriboon, M., Kosawat, K., & Supnithi, T. (2020). *The Annotation Guideline of LST20 Corpus*. ArXiv. <https://doi.org/10.48550/ARXIV.2008.05055>.
- Chormai, P., Prasertsom, P., & Rutherford, A. (2019). *AttaCut: A Fast and Accurate Neural Thai Word Segmenter*. ArXiv. <https://doi.org/10.48550/ARXIV.1911.07056>.
- Conneau, A., Khandelwal, K., Goyal, N., Chaudhary, V., Wenzek, G., Guzmán, F., Grave, E., Ott, M., Zettlemoyer, L., & Stoyanov, V. (2020). Unsupervised Cross-lingual Representation Learning at Scale. *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 8440–8451. <https://doi.org/10.18653/v1/2020.acl-main.747>.
- De Vries, W., Wieling, M., & Nissim, M. (2022). Make the Best of Cross-lingual Transfer: Evidence from POS Tagging with over 100 Languages. *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics*, 7676–7685. <https://doi.org/10.18653/v1/2022.acl-long.529>.
- Haruechaiyasak, C., Kongyoung, S., & Dailey, M. (2008). A comparative study on Thai word segmentation approaches. *2008 5th International Conference on Electrical Engineering/Electronics, Computer, Telecommunications and Information Technology*, 1, 125–128. <https://doi.org/10.1109/ECTICON.2008.4600388>.
- Kittinaradorn, R., Chaovavanich, K., Achakulvisut, T., Srithaworn, K., Chormai, P., Kaewkasi, C., Ruangrong, T., & Oparad, K. (2019). *DeepCut: A Thai Word Tokenization Library using Deep Neural Network*. (v1.0). Zenodo. <https://doi.org/10.5281/zenodo.3457707>.
- National Electronics and Computer Technology Center. (2022). LST20 Corpus. <https://opend-portal.nectec.or.th/it/dataset/lst20-corpus>. (In Thai)
- Noyunsan, C., Haruechaiyasak, C., Poltree, S., & Saikaew, K.R. (2014). A Multi-Aspect Comparison and Evaluation on Thai Word Segmentation Programs. *Joint International Conference of Semantic Technology*. https://ceur-ws.org/Vol-1312/jist2014pd_paper6.pdf.
- Poowarawan, Y. (1986). Dictionary-based Thai Syllable Separation. *In Proceedings of the Ninth Electronics Engineering Conference*, 409–418.
- Ruenlek, K., & Damrongkamoltip, K. (2023). *A Mockup Dataset for Word Segmentation Based on the LST20 Dataset*. <https://github.com/Fairpart/A-mockup-dataset-for-word-segmentation-based-on-the-LST20-DATASET>.
- Theeramunkong, T., Sornlertlamvanich, V., Tanhermhong, T., & Chinnan, W. (2000). Character Cluster based Thai Information Retrieval. *Proceedings of the Fifth International Workshop on on Information Retrieval with Asian Languages*, 75–80. <https://doi.org/10.1145/355214.355225>.
- Virach, S. (1993). Word Segmentation for Thai in Machine Translation System. *Machine Translation*. 50-56. (In Thai)

Land Price Prediction with Machine Learning in Mueang Khon Kaen District

การทำนายราคาที่ดินด้วยการเรียนรู้ของเครื่องในอำเภอเมืองขอนแก่น

Yosita Sriwuttisap¹, Dusita Sungklinhom², Sakpod Tongleamnak³, and Thanaphon Tangchoopong^{4,*}

โยษิตา ศรีวุฒิสัพ¹, ดุสิตา สังข์กลิ่นหอม², ศักดิ์พจน์ ทองเลี่ยมนาค³, และ ธนพล ตั้งชูพงศ์^{4,*}

Received: 20 March 2024;

Revised: 18 May 2024;

Accepted: 22 May 2024;

Published: 15 August 2024;

Abstract

This research aims to develop a model of land price assessment and study of factors influencing land prices with machine learning used as a guideline for determining the estimated price to be close to the actual purchase price in Mueang Khon Kaen District from land price information traded on websites in enforcement department of 193 locations, and land price assessment information in land department of 1,500 locations in this study the factors involved include appraisal value, property type, land size, distance, and the average appraisal value of five nearby plots of land. The models used for analysis are Regression Tree, Random Forest, Gradient Boosted Trees, and Linear Regression. Which, the land price assessment from the case study by model measurement in MAE, RMSE, R-squared, Grid Search, and Cross-validation to selected model parameters and evaluate their performance. The results found that the model with the best predictive performance is Gradient Boosted Trees in R-squared at the highest of 0.80, MAE, and RMSE at the lowest of 7929.40, and 15281.33, respectively. Feature Importance in the locations with the most influence on prediction, followed by area size, average appraised value from five nearby locations, and property type.

Keywords: Land Price Assessment, Machine Learning, Important Factors.

¹ Student, Department of Computing Science, College of Computing, Khon Kaen University, Khon Kaen 40002, Thailand; นักศึกษา สาขาวิชาวิทยาการคอมพิวเตอร์ วิทยาลัยการคอมพิวเตอร์ มหาวิทยาลัยขอนแก่น จังหวัดขอนแก่น 40002 ประเทศไทย; Email: yosita.sr@kkumail.com

² Student, Department of Computing Science, College of Computing, Khon Kaen University, Khon Kaen 40002, Thailand; นักศึกษา สาขาวิชาวิทยาการคอมพิวเตอร์ วิทยาลัยการคอมพิวเตอร์ มหาวิทยาลัยขอนแก่น จังหวัดขอนแก่น 40002 ประเทศไทย; Email: dusita.su@kkumail.com

³ Lecturer, Dr., Department of Computing Science, College of Computing, Khon Kaen University, Khon Kaen 40002, Thailand; อาจารย์ ดร. สาขาวิชาวิทยาการคอมพิวเตอร์ วิทยาลัยการคอมพิวเตอร์ มหาวิทยาลัยขอนแก่น จังหวัดขอนแก่น 40002 ประเทศไทย; Email: sakpod@kku.ac.th

^{4,*} Lecturer, Department of Computing Science, College of Computing, Khon Kaen University, Khon Kaen 40002, Thailand; อาจารย์ สาขาวิชาวิทยาการคอมพิวเตอร์ วิทยาลัยการคอมพิวเตอร์ มหาวิทยาลัยขอนแก่น จังหวัดขอนแก่น 40002 ประเทศไทย; Email: thanaphon@kku.ac.th

*Corresponding authors: Thanaphon Tangchoopong (thanaphon@kku.ac.th)



บทคัดย่อ

การวิจัยนี้มีวัตถุประสงค์เพื่อพัฒนาโมเดลการประเมินราคาที่ดิน และศึกษาปัจจัยที่มีอิทธิพลต่อราคาที่ดิน ด้วยการเรียนรู้ของเครื่อง ที่ใช้เป็นแนวทางของการกำหนดราคาประเมินให้ได้ใกล้เคียงกับราคาซื้อขายจริงในอำเภอเมืองขอนแก่น จากข้อมูลราคาที่ดินที่มีการซื้อขายในเว็บไซต์กรมบังคับคดี 193 แห่ง และข้อมูลราคาประเมินของกรมที่ดิน 1500 แห่ง ในการศึกษาปัจจัยที่เกี่ยวข้องได้แก่ ราคาประเมิน ประเภททรัพย์สิน ขนาดพื้นที่ ระยะทาง และราคาประเมินเฉลี่ยจากที่ดินห้าผืนใกล้เคียง โมเดลที่ใช้ในการวิเคราะห์ ได้แก่ ต้นไม้ตัดสินใจแบบถดถอย แรดอมฟอเรสต์ เกรเดียนบูสท์ และการถดถอยเชิงเส้น ซึ่งการประเมินราคาที่ดินจากกรณีศึกษาโดยการวัดผลโมเดลในค่าความคลาดเคลื่อนเฉลี่ยสัมบูรณ์ ค่าความคลาดเคลื่อนเฉลี่ยรากที่สอง ค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวน การค้นหาแบบกริด และการประเมินประสิทธิภาพแบบไขว้พบเพื่อคัดเลือกพารามิเตอร์ของโมเดล และประเมินประสิทธิภาพ ผลวิจัยพบว่า โมเดลที่มีผลการทำนายข้อมูลที่ดีที่สุดคือ เกรเดียนบูสท์ ที่มีค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวนสูงสุดเท่ากับ 0.80 ค่าความคลาดเคลื่อนเฉลี่ยสัมบูรณ์ และค่าความคลาดเคลื่อนเฉลี่ยรากที่สองต่ำสุด เท่ากับ 7929.40 และ 15281.33 ตามลำดับ ความสำคัญของคุณลักษณะในกลุ่มสถานที่ที่มีผลต่อการทำนายมากที่สุด รองลงมาคือขนาดพื้นที่ ราคาประเมินเฉลี่ยจากห้าตำแหน่งใกล้เคียง และประเภททรัพย์สิน

คำสำคัญ: ประเมินราคาที่ดิน, การเรียนรู้ของเครื่อง, ปัจจัยสำคัญ

1. บทนำ (Introduction)

พื้นที่ในเขตอำเภอเมือง จังหวัดขอนแก่น มีการพัฒนาและเติบโตอย่างต่อเนื่อง เป็นศูนย์กลางในการขยายธุรกิจในภูมิภาคตะวันออกเฉียงเหนือรวมถึงธุรกิจอสังหาริมทรัพย์ ซึ่งส่งผลให้ความต้องการที่ดินเพื่อการพัฒนาเพิ่มสูงขึ้น ราคาที่ดินคือมูลค่าปัจจุบันของผลประโยชน์จากการใช้ที่ดิน และเกี่ยวข้องกับธุรกรรมทางการเงินและเศรษฐกิจ การกำหนดราคาที่ดินในประเทศไทยดำเนินการโดยทั้งหน่วยงานภาครัฐและเอกชน โดยกรมธนารักษ์ กระทรวงการคลัง เป็นหน่วยงานหลักที่รับผิดชอบการประเมินและจัดทำบัญชีราคาที่ดินเพื่อใช้เป็นฐานในการจัดเก็บภาษีและค่าธรรมเนียมต่าง ๆ อย่างไรก็ตาม ราคาประเมินที่ดินมักต่ำกว่าราคาตลาดจริง (Kumpu & Piyathamronchai, 2024) ด้วยเหตุนี้การพัฒนาโมเดลในการประเมินราคาที่ดินจึงมีความสำคัญอย่างยิ่งเพื่อช่วยภาครัฐ และให้ผู้สนใจลงทุนในที่ดินมีข้อมูลที่ดีและแม่นยำมากยิ่งขึ้น

งานวิจัยนี้มุ่งเน้นการพัฒนาแบบจำลองเพื่อประเมินราคาที่ดินในเขตอำเภอเมือง จังหวัดขอนแก่นโดยใช้ข้อมูลที่มีความเกี่ยวข้อง โดยแบ่งออกเป็นสามกลุ่มปัจจัย ได้แก่ ปัจจัยด้านราคาประเมิน ปัจจัยด้านขนาดและปัจจัยทำเลที่ตั้ง และปัจจัยด้านประเภททรัพย์สิน (Soltani et al., 2022) ดังแสดงในกรอบแนวคิดในการวิจัย Figure 1. ซึ่งข้อมูลที่ได้ถูกนำไปพัฒนาโมเดล และอาศัยการศึกษาปัจจัยที่ส่งผลต่อการทำนาย เพื่อใช้ในการทดลองปรับปรุงประสิทธิภาพของโมเดลผ่านกระบวนการปรับตัวแปรต้น (Feature Selection) ดำเนินการเลือกใช้เทคนิคการคัดเลือกโมเดล (Model Selection) จากโมเดลต้นไม้ตัดสินใจแบบถดถอย แรดอมฟอเรสต์ เกรเดียนบูสท์ และสมการถดถอยเชิงเส้น (Soltani et al., 2022; Burnham & Anderson, 2002) สำหรับประเมินราคาที่ดินกรณีศึกษาเขตอำเภอเมือง จังหวัดขอนแก่นและวัดผลโมเดลโดย ค่าความคลาดเคลื่อนเฉลี่ยสัมบูรณ์ ค่าความคลาดเคลื่อนเฉลี่ยรากที่สอง และค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวน

2. วัตถุประสงค์งานวิจัย (Research Objectives)

1. พัฒนาโมเดลการประเมินราคาที่ดินในเขตอำเภอเมืองจังหวัดขอนแก่น
2. เพื่อศึกษาปัจจัยที่มีอิทธิพลต่อราคาที่ดินโดยใช้การเรียนรู้ของเครื่อง

3. กรอบแนวคิดงานวิจัย (Conceptual Framework)

งานวิจัยนี้กำหนดกรอบแนวคิดในการทำงานวิจัยและดำเนินการตามแนวคิดทฤษฎี ดัง Figure 1.

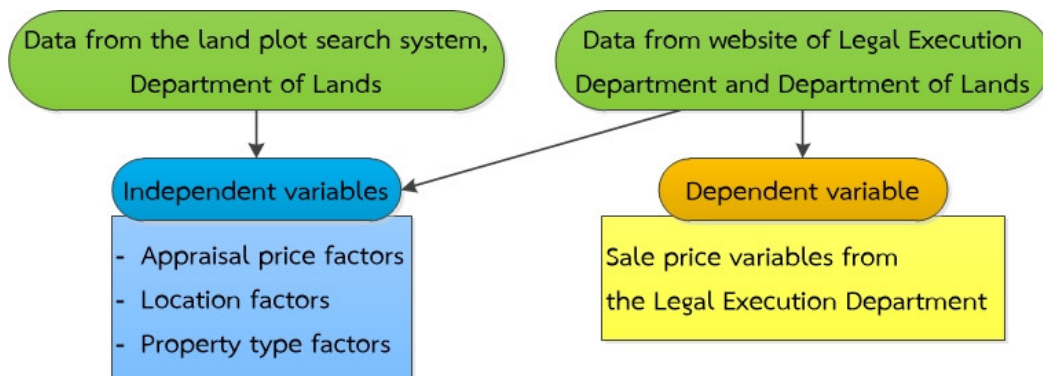


Figure 1. Conceptual Framework

4. การทบทวนวรรณกรรมและทฤษฎีที่เกี่ยวข้อง (Literature Review)

Soltani et al. (2022) ได้นำเสนอโมเดลเรียนรู้ของเครื่องทั้งหมด 4 รูปแบบเพื่อศึกษาผลของลักษณะต่าง ๆ เช่น คุณลักษณะของทรัพย์สินและคุณภาพของย่านบริเวณที่อยู่ต่อการแปรผันของราคาที่พักอาศัยในพื้นที่ภูมิศาสตร์ที่แตกต่าง โดยใช้ชุดข้อมูลราคาที่พักอาศัยใน Metropolitan Adelaide, ออสเตรเลีย ในระยะเวลา 32 ปี (1984 - 2016) การวิจัยนี้ขึ้นอยู่กับข้อมูลการทำธุรกรรมการขายที่มีจำนวนมากถึง 428,000 รายการและตัวแปรอธิบาย 38 ตัวแปร พบว่าโมเดลที่ใช้ต้นไม้ตัดสินใจแบบถดถอย มีประสิทธิภาพมากกว่าโมเดลเชิงเส้น นอกจากนี้ เทคนิคการเรียนรู้แบบกลุ่ม เช่น แรนดอมฟอเรสต์ และเกรเดียนบูสท์ มีประสิทธิภาพมากกว่าในการทำนายราคาที่พักอาศัยในอนาคต โดยได้เพิ่มตัวแปรช่วงเวลาและพื้นที่ (Spatio-Temporal Lag) เพื่อเพิ่มความแม่นยำในการทำนายของโมเดล การวิจัยนี้ชี้ให้เห็นว่าตัวแปรช่วงเวลาและพื้นที่ (หรือตัวบ่งชี้ของพื้นที่ เวลาที่คล้ายกัน) เป็นปัจจัยที่สำคัญในการทำนายราคาที่พักอาศัยในอนาคต

Kumpu & Piyathamronchai (2024) ได้พัฒนาแบบจำลองการประเมินราคาที่ดินในอำเภอเมืองเชียงใหม่ จังหวัดเชียงใหม่ โดยใช้ข้อมูลการซื้อขายและจดทะเบียนสิทธิและนิติกรรมจากสำนักงานที่ดินจังหวัดเชียงใหม่ระหว่างเดือนพฤษภาคม พ.ศ. 2560 ถึง พฤศจิกายน พ.ศ. 2565 จำนวน 2,821 รายการ และใช้แบบสอบถามออนไลน์ (Google Forms) จำนวน 300 ราย การวิเคราะห์ข้อมูลใช้สถิติเชิงพรรณนาและสถิติเชิงอนุมาน รวมถึงการถดถอยเชิงเส้นพหุคูณแบบเป็นขั้นตอน (Stepwise Regression) พบว่าปัจจัยที่มีอิทธิพลต่อราคาประเมินที่ดินในทิศทางบวก ได้แก่ มูลค่าถนน ความกว้างของแปลงที่ดิน ระยะห่างจากแหล่งมลภาวะ และความครบถ้วนของสาธารณูปโภค ส่วนปัจจัยที่มีอิทธิพลในทิศทางลบ ได้แก่ เนื้อที่แปลงที่ดิน ผลการศึกษาพบว่าแบบจำลองสามารถอธิบายราคาที่ดินได้เพียงร้อยละ 56.9 ซึ่งชี้ให้เห็นว่ายังมีปัจจัยอื่น ๆ ที่ไม่ได้ถูกนำมาพิจารณา ซึ่งควรมีการศึกษาปัจจัยเพิ่มเติม เช่น ปัจจัยทางเศรษฐกิจและการเก็งกำไร เพื่อเพิ่มความแม่นยำของแบบจำลอง นอกจากนี้ยังมีข้อจำกัดเรื่องช่วงเวลาของข้อมูลที่อาจส่งผลต่อความแม่นยำในการวิเคราะห์และประมวลผล

4.1 แฮเวอร์ซีน (Haversine)

การคำนวณระยะทางวงกลม ระหว่างตัวอย่างในเซต X และ Y บนพื้นผิวของทรงกลม พิกัดแรกของแต่ละจุดถือเป็นละติจูด และพิกัดที่สองถือเป็นลองจิจูด และมีหน่วยเป็นเรเดียน (Pedregosa et al., 2012)

$$D_{(x,y)} = 2\arcsin \left[\sqrt{\sin^2 \left(\frac{x_1-y_1}{2} \right) + \cos(x_1) \cos(y_1) \sin^2 \left(\frac{x_2-y_2}{2} \right)} \right] \quad (1)$$

4.2 อิทธิพลของคุณลักษณะ

เป็นเทคนิคที่ใช้ในการระบุว่าแต่ละคุณลักษณะ (Feature) มีความสำคัญต่อผลลัพธ์ของโมเดลมากน้อยแค่ไหน การประเมินความสำคัญของคุณลักษณะช่วยให้เข้าใจว่าโมเดลให้ความสำคัญกับข้อมูลส่วนใดมากที่สุด ซึ่งสามารถใช้ในการปรับปรุงโมเดลหรือการตัดสินใจทางธุรกิจ ตัวอย่างเทคนิคที่ใช้ในการประเมินความสำคัญของคุณลักษณะได้แก่ การใช้ค่าสัมประสิทธิ์ (Coefficient) ของโมเดลเชิงเส้น (Linear) การใช้ค่าความสำคัญของคุณลักษณะในต้นไม้ตัดสินใจ หรือการใช้ค่า SHAP values (Vanderplas, 2017; Lundberg & Lee, 2017)

4.3 ปัญหาสมการถดถอย (Regression Problem)

ปัญหาการถดถอย (Na Bangchang, 2011) เป็นกระบวนการในการสร้างแบบจำลองทางสถิติหรือการเรียนรู้ของเครื่องเพื่อทำนายค่าตัวแปรตาม (Dependent Variable) ที่เป็นค่าต่อเนื่อง (Continuous Variable) จากค่าตัวแปรอิสระ (Independent Variables) ที่เป็นตัวแปรเชิงอธิบายหรือคุณลักษณะต่าง ๆ (Features) โดยมีเป้าหมายในการหาเส้นหรือฟังก์ชันที่สามารถอธิบายความสัมพันธ์ระหว่างตัวแปรเหล่านี้ได้อย่างดีที่สุด ซึ่งแบบจำลองที่ใช้แก้ปัญหาประเภทนี้ ได้แก่

1. การวิเคราะห์การถดถอยเชิงเส้น

เป็นการนำเอาข้อมูลหรือตัวแปรมาหาความสัมพันธ์กันโดยเมื่อมีตัวแปรอิสระเพียงตัวเดียวเรียกว่า การถดถอยอย่างง่าย (Simple Linear Regression) เมื่อมีตัวแปรอิสระมากกว่า 1 ตัวแปรเรียกว่าการวิเคราะห์การถดถอยเชิงพหุคูณ (Multiple Linear Regression) (Na Bangchang, 2011) โดยมีรูปแบบสมการในการวิเคราะห์ดังนี้

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n + e \quad (2)$$

กำหนดให้

Y	คือ ตัวแปรเกณฑ์
$X_1 X_2 \dots X_n$	คือ ค่าของตัวแปรอิสระแต่ละตัว
β_0	คือ ส่วนตัดแกน Y
$\beta_2 \dots \beta_n$	คือ ค่าสัมประสิทธิ์ความถดถอย
e	คือ ค่าความคลาดเคลื่อน (Error Residual)

2. ต้นไม้ตัดสินใจแบบถดถอย (Regression Trees)

เป็นการสร้างเงื่อนไข If-else ขึ้นมาจากข้อมูลในตัวแปรเพื่อแบ่งข้อมูลเป็นกลุ่มใหม่ที่อธิบายเป้าหมาย (Target) ได้ดีที่สุด ต้นไม้ตัดสินใจเราใช้ฟังก์ชันวัตถุประสงค์ (Objective Function) ที่เหมาะสมเพื่อกำหนดแยก (Split) ในแต่ละตัวแปร โดยแต่ละประเภทของ ต้นไม้ตัดสินใจมีฟังก์ชันวัตถุประสงค์ต่างกัน เช่น จีนิอิมพริที (Gini Impurity) หรือเอนโทรปี (Entropy) สำหรับต้นไม้ตัดสินใจแบบจำแนก (Classification Tree) และค่าผลรวมของผลต่างกำลังสอง (Residual

sum of squares : RSS) สำหรับต้นไม้ตัดสินใจแบบถดถอย (Regression Tree) ซึ่งจะช่วยในการแบ่งข้อมูลให้อยู่ในกลุ่มที่เหมาะสมสำหรับการอธิบายตัวแปรเป้าหมาย (Target Variable) (Vanderplas, 2017)

$$e_i = y_i - \hat{y}_i \quad (3)$$

$$RSS = e_1^2 + e_2^2 + e_3^2 + \dots + e_n^2 \quad (4)$$

Residual (Error Term, e_i) คือ ค่าความคลาดเคลื่อนหรือค่าผิดพลาด (Error) ระหว่าง y ทุก ๆ จุดในข้อมูล กับ $\hat{y}$ ที่ได้มาจากการประมาณค่าการทำนาย (Prediction) ขึ้นมาการคำนวณค่าความคลาดเคลื่อนของข้อมูลตัวที่ i

3. แรนดอมฟอเรสต์ (Random Forest)

เป็นแบบจำลองที่ถูกพัฒนาขึ้นจากต้นไม้ตัดสินใจ โดยที่แรนดอมฟอเรสต์จะทำการเพิ่มจำนวนต้นไม้เป็นหลาย ๆ ต้นโดยใช้กระบวนการสุ่มข้อมูล (Bootstrap Sampling) และสุ่มคุณลักษณะ ในแต่ละโหนดของต้นไม้เพื่อลดความเกี่ยวข้องระหว่างต้นไม้แต่ละต้น และทำนายด้วยการโหวตหรือหาค่าเฉลี่ย (Vanderplas, 2017)

4. เกรเดียนบูสทรี (Gradient boosted Trees)

เป็นแบบจำลองที่ถูกพัฒนาขึ้นจากต้นไม้ตัดสินใจโดยสร้างต้นไม้ทีละต้นโดยทำนายค่าผลลัพธ์และคำนวณค่าความผิดพลาด ระหว่างค่าทำนายกับค่าจริงจากนั้นในทุกรอบจะทำการสร้างต้นไม้เพิ่มเติมเพื่อลดความผิดพลาดที่เหลือ (Vanderplas, 2017)

4.4 การปรับแต่งไฮเปอร์พารามิเตอร์ และการคัดเลือกโมเดล

1. การปรับแต่งไฮเปอร์พารามิเตอร์ (Hyper-Parameter Tuning)

กระบวนการการปรับแต่งค่าพารามิเตอร์ที่ใช้ในการฝึกโมเดลเครื่องจักรอัลกอริทึมเพื่อให้โมเดลมีประสิทธิภาพสูงสุด หรือให้ผลลัพธ์ที่ดีที่สุดที่เป็นไปได้สำหรับงานที่กำลังดำเนินการอยู่ (Vanderplas, 2017)

2. รูปแบบของการประเมินผล (Model Evaluation)

เป็นกระบวนการที่ใช้เพื่อวัดประสิทธิภาพของโมเดลที่ได้ฝึกและทดสอบด้วยข้อมูลที่ไม่เคยเห็นในกระบวนการฝึก (Unseen Data) เพื่อประเมินความสามารถในการทำนายหรือจำแนกข้อมูลใหม่ที่ไม่เคยเห็นมาก่อน กระบวนการนี้เป็นส่วนสำคัญของการพัฒนาและใช้งานแบบจำลอง (Vanderplas, 2017)

3. การทดสอบแบบไขว้ทบ (Cross Validation)

เป็นเทคนิคที่ใช้ในการประเมินความสามารถของโมเดลและช่วยในการป้องกันปัญหาโอเวอร์ฟิต (Overfitting) โดยทั่วไปจะใช้วิธีการแบ่งข้อมูลออกเป็นหลายส่วน (Fold) ซึ่งเรียกว่าการทดสอบแบบไขว้ทบแบบ K ส่วน (K-Fold Cross-Validation) (Kohavi, 1995) วิธีการนี้จะทำการแบ่งข้อมูลเป็น K ส่วน จากนั้นทำการฝึกโมเดลด้วยข้อมูล K-1 ส่วน และใช้ส่วนที่เหลือในการทดสอบโมเดล ทำแบบนี้ K รอบ เพื่อให้มั่นใจว่าโมเดลได้รับการทดสอบกับข้อมูลทั้งหมด กระบวนการนี้ช่วยให้เราสามารถประเมินประสิทธิภาพของโมเดลได้อย่างแม่นยำยิ่งขึ้น

4. การคัดเลือกโมเดล

เป็นกระบวนการเลือกโมเดลที่ดีที่สุดจากชุดของโมเดลที่ใช้ทำนายหรือจัดกลุ่มข้อมูล กระบวนการนี้รวมถึงการเปรียบเทียบประสิทธิภาพของโมเดลต่างๆ โดยใช้ metrics เช่น ความแม่นยำ (Accuracy) ค่า F1 score ค่าความคลาดเคลื่อนกำลังสองเฉลี่ย (Mean Squared Error: MSE) เป็นต้น นอกจากนี้ยังรวมถึงการเลือกค่าพารามิเตอร์ที่เหมาะสม

สำหรับโมเดลเพื่อให้ได้ผลลัพธ์ที่ดีที่สุด รวมไปถึงการใช้การทดสอบแบบไขว้ทับก็เป็นส่วนสำคัญในการเลือกโมเดลที่ดีที่สุด (Burnham & Anderson, 2002)

5. การค้นหาพารามิเตอร์แบบกริด (Grid Search for Hyperparameter Tuning)

การค้นหาพารามิเตอร์แบบกริด (Bergstra & Bengio, 2012) เป็นวิธีการปรับแต่งค่าพารามิเตอร์ของโมเดล โดยการค้นหาชุดค่าที่เหมาะสมที่สุดในพื้นที่ของพารามิเตอร์ที่กำหนดไว้ล่วงหน้า วิธีการนี้จะทำการทดสอบค่าพารามิเตอร์ต่างๆ ที่เป็นไปได้ทั้งหมดอย่างเป็นระบบและครอบคลุมทุกค่าในกริดที่กำหนด จากนั้นจะประเมินประสิทธิภาพของแต่ละชุดค่าพารามิเตอร์ด้วยการใช้ การทดสอบแบบไขว้ทับเพื่อหาโมเดลที่มีประสิทธิภาพสูงสุด กระบวนการนี้ช่วยเพิ่มความแม่นยำและความสามารถในการทำนายของโมเดล

6. กระบวนการคัดเลือกคุณลักษณะ

กระบวนการเลือกคุณลักษณะมีความสำคัญและเป็นประโยชน์ในการสร้างโมเดลจากชุดคุณลักษณะที่มีอยู่หลายๆ ตัว โดยมีเป้าหมายเพื่อเพิ่มประสิทธิภาพของโมเดล ลดความซับซ้อน และลดความเสี่ยงของการโอเวอร์ฟิต กระบวนการนี้ช่วยให้โมเดลสามารถทำงานได้เร็วขึ้นและแม่นยำมากขึ้นมีหลายวิธีที่ใช้ในการเลือกคุณลักษณะเช่น

ฟิวเจอร์เมธอด (Filter Methods) วิธีการนี้จะเลือกคุณลักษณะโดยพิจารณาจากคุณสมบัติทางสถิติ เช่น การใช้ค่าสหสัมพันธ์ (Correlation) ค่าสถิติไคสแควร์ (Chi-Square Test) หรือ สารสนเทศร่วม (Mutual Information) โดยวิธีนี้ไม่ต้องใช้โมเดลในการเลือกคุณลักษณะ

วิธีแรปปอร์ (Wrapper Methods) วิธีการนี้จะเลือกคุณลักษณะโดยใช้โมเดลเป็นส่วนหนึ่งของกระบวนการเลือก เช่น การกำจัดคุณลักษณะด้วยกระบวนการ

การเวียนเกิด (Recursive Feature Elimination: RFE) ที่ทำการลบคุณลักษณะที่มีความสำคัญน้อยที่สุดทีละตัว จนกว่าจะได้ชุดคุณลักษณะที่ดีที่สุด

วิธีฝังตัว (Embedded Methods) วิธีการนี้จะเลือกคุณลักษณะในขณะที่สร้างโมเดล เช่น การใช้สมการถดถอยแบบลาสโซ่ (Lasso Regression) ที่สามารถทำให้ค่าของสัมประสิทธิ์ของคุณลักษณะที่ไม่สำคัญเป็นศูนย์ (Vanderplas, 2017) หรือการใช้คุณลักษณะสำคัญของ ของกระบวนการต้นไม้ตัดสินใจแบบถดถอย และแรนดอมฟอเรสต์

การเลือกคุณลักษณะที่เหมาะสม (Guyon & Elisseeff, 2003; Chandrashekar & Sahin, 2014; Choompol, 2019) ช่วยเพิ่มประสิทธิภาพของโมเดล ลดเวลาในการประมวลผล และทำให้การวิเคราะห์ผลลัพธ์ของโมเดลง่ายขึ้น

5. วิธีดำเนินงานวิจัย (Research Methodology)

งานวิจัยนี้ได้ทำการรวบรวมข้อมูล (Data Collection) ข้อมูลรายงานผลการขายทอดตลาด ข้อมูลค้นหาทรัพย์สิน ประกาศขายทอดตลาด และข้อมูลระบบคั่นหารูปแปลงที่ดิน ระบบคั่นหารูปแปลงที่ดินของกรมที่ดินจากนั้นทำการสร้างคุณลักษณะใหม่ (Feature Construction) ตัวแปรระยะทางจากสถานที่สำคัญ 14 สถานที่หน่วยเมตร (แฮเวอร์ซัน) ตัวแปรราคาประเมินจากค่าเฉลี่ย 5 สถานที่ใกล้เคียง ตัวแปรขนาดพื้นที่หน่วยตารางวาและประเภททรัพย์สิน

ต่อมาทำการวิเคราะห์ข้อมูลเชิงการตรวจสอบและกระบวนการเตรียมข้อมูล (EDA & Data Pre Processing) โดยทำการนำเสนอแผนภาพข้อมูล (Data Visualization) การจัดการข้อมูลที่หายไป (Missing Value Handling) การตรวจสอบค่าผิดปกติ (Outlier Detection) และการเข้ารหัสข้อมูลแบบหมวดหมู่ (Encoding Categorical Data) จากนั้นนำข้อมูลมาทำการเรียนรู้ของเครื่องโดยใช้ 4 โมเดลดังนี้ ต้นไม้ตัดสินใจ แรนดอมฟอเรสต์ เกรเดียนบูสท์และการถดถอยเชิงเส้นโดยมีการใช้การค้นหาแบบกริด เป็นการค้นหาค่าพารามิเตอร์โดยการกำหนดค่าที่เป็นไปได้ทั้งหมดของแต่ละพารามิเตอร์ที่สนใจในรูปแบบของรายการที่กำหนดล่วงหน้า และสุดท้ายรูปแบบของการประเมินผล (Models Evaluation)

โดยใช้ค่าความคลาดเคลื่อนเฉลี่ยสัมบูรณ์ ค่าความคลาดเคลื่อนเฉลี่ยรากที่สอง และค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวนในการวัดและทดสอบประสิทธิภาพของแบบจำลอง

5.1 การเก็บรวบรวมข้อมูล (Data Collection)

1. พื้นที่กรณีศึกษาราคาขายจริง 193 แห่ง

ดำเนินการเก็บรวบรวมข้อมูลจากเว็บไซต์กรมบังคับคดีและกรมที่ดินโดยที่ดินจากกรมบังคับคดีในพื้นที่เขตอำเภอเมืองจังหวัดขอนแก่นระหว่างปี 2560 ถึง 2566 ข้อมูลใน 1 ระเบียบแสดงข้อมูลตัวอย่างตาม Table 1. โดยข้อจำกัดของชุดข้อมูลจากเว็บไซต์กรมบังคับคดีเมื่อทำการขายเสร็จสิ้นข้อมูลจะถูกลบออกไปบางส่วน ซึ่งการเก็บข้อมูลย้อนหลังอาจทำให้ข้อมูลการจ้างนองติดไปของที่ดินอาจสูญหายโดยในงานวิจัยนี้ไม่ได้นำข้อมูลการติดจ้างนองมาเป็นตัวแปรต้น

Table 1. The case study of land detail of 193 locations.

Data	Example
Case number	ผบ.1003
Property type	Land with Structures
Rai (ไร่)	-
Ngan (งาน)	-
Square Wah (ตารางวา)	73.4
Appraised Value	3,196,000.00
Sub-district	Banped
District	Mueang Khon Kaen
Province	Khon Kaen
Title Deed Land	254492
Achievable Selling Price/ Highest Bid	3,350,000
Cadastral Map	5541 I 6414-05 (1000)

Table 2. The color setting of case study of land detail of 193 locations.

Group	The selling price range per square wah		Average selling price	Color	Number
	lower bound	upper bound			
1	150000	300000	205675	White	7
2	100000	149999	131298	Red	4
3	50000	99999	63107	Blue	9
4	10000	49999	25579	Green	66
5	5000	9999	7114	Orange	21
6	1000	4999	2425	Purple	54
7	300	999	589	Pink	28
8	1	299	180	Brown	4



จาก Table 2. กำหนดสีและสัญลักษณ์สำหรับกลุ่มพล็อตกระจายเพื่อดูการแบ่งกลุ่มตามราคาขายต่อตารางวา โดยกำหนดกลุ่มของราคาประเมินเป็น 8 กลุ่มตามตัวเลขเริ่มต้นและตัวเลขสิ้นสุดของราคา

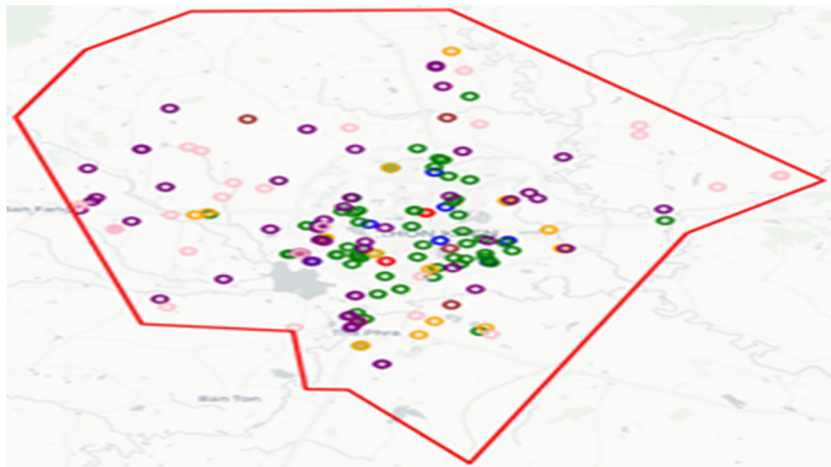


Figure 2. ตำแหน่งการกระจายตัว ของข้อมูลราคาขายต่อตารางวา ซึ่งเป็นข้อมูลราคาขายจริงจากเว็บไซต์กรมบังคับคดี

จาก Figure 2. แสดงข้อมูลตำแหน่งบนแผนที่เพื่อแสดงข้อมูลที่คัดแยกตามกลุ่มราคาขายต่อตารางวาที่กำหนดกลุ่มของราคาประเมินเป็น 8 กลุ่มตามตัวเลขเริ่มต้นและตัวเลขสิ้นสุดของราคา และเส้นสีแดงแสดงถึงกรอบบริเวณอำเภอเมืองขอนแก่น

2. พื้นที่กรณีศึกษาราคาประเมิน 1500 แห่ง

ผู้วิจัยได้ทำการสุ่มข้อมูลราคาประเมินที่ดินตามกลุ่มราคาที่แสดงใน Table 3. จากระบบค้นหาแปลงที่ดิน (Landsmaps) กรณีที่ดินจำนวน 1500 แห่ง ซึ่งใน 1 ระเบียบ ประกอบด้วย

(1) ละติจูดและลองจิจูดจากที่ดินโดยการสุ่มตัวอย่างแบบชั้นภูมิ (Stratified Sampling) โดยกำหนดกลุ่มของราคาประเมินเป็น 6 กลุ่มตามตัวเลขเริ่มต้นและตัวเลขสิ้นสุด (Range) ของราคา

(2) ข้อมูลราคาประเมินบาทต่อตารางวาจาก Table 3. กำหนดสีและสัญลักษณ์สำหรับกลุ่มพล็อตกระจายเพื่อดูการแบ่งกลุ่มตามราคาต่อตารางวาจากกรณีศึกษาที่ดิน 1,500 แห่งที่มีการสุ่มตัวอย่างแบบชั้นภูมิโดยกำหนดกลุ่มของราคาประเมินเป็น 6 กลุ่มตามตัวเลขเริ่มต้นและตัวเลขสิ้นสุดของราคา

Table 3. The color setting of case study of land detail of 1,500 locations.

Group	The selling price range per square wah		Average selling price	Color	Number
	lower bound	upper bound			
1	100000	150000	122736	Pink	250
2	50000	99999	63088	Red	250
3	10000	49999	23168	Blue	250
4	5000	9999	6940	Green	250
5	1000	4999	2086	Orange	250
6	300	999	579	Purple	250

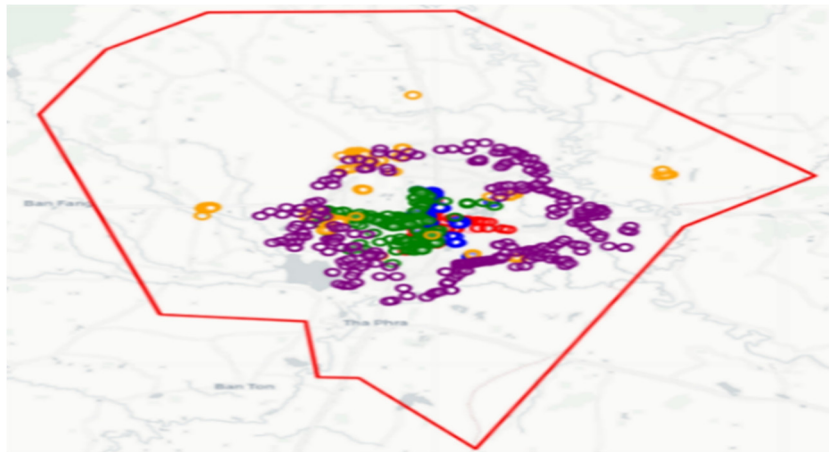


Figure 3. Scatter plot to observe clustering based on the price per square wah (appraisal price) in a case study of 1,500 locations.

จาก Figure 3. แสดงข้อมูลตำแหน่งบนแผนที่เพื่อแสดงข้อมูลที่คัดแยกตามกลุ่มราคาต่อตารางวาที่กำหนดกลุ่มของราคาประเมินเป็น 6 กลุ่ม และเส้นสีแดงล้อมบริเวณอำเภอเมืองขอนแก่น

5.2 การสร้างคุณลักษณะ (Feature Construction)

1. สกัดคุณลักษณะสำคัญจากระยะทาง โดยใช้แฮเวอร์ซีน ซึ่งเป็นระยะทางเชิงมุมระหว่างจุดสองจุดบนพื้นผิวของทรงกลมละติจูดลองติจูด จากที่ดินที่สนใจไปหาสถานที่สำคัญ 14 สถานที่ โดยแบ่งเป็น สวนสาธารณะ โรงแรม ห้างสรรพสินค้า พิพิธภัณฑ์ สนามบิน โรงพยาบาล สถาบันการศึกษา ตลาด และศาสนสถานใช้หน่วยระยะทางเป็นกิโลเมตร
2. สกัดราคาประเมินจากค่าเฉลี่ย 5 สถานที่ใกล้เคียงผู้จัดทำคำนวณระยะทางระหว่างพิกัดที่ดินของกรณีศึกษาที่ดิน 1,500 แห่ง ในอำเภอเมือง จังหวัดขอนแก่นแต่ละแถวกับพิกัดเป้าหมายของกรณีศึกษาที่ดิน 193 แห่ง ในอำเภอเมือง จังหวัดขอนแก่น เรียงลำดับตามระยะทางเพื่อหาตำแหน่งที่ใกล้ที่สุด 5 ตำแหน่งแล้วนำ 5 ตำแหน่งที่ทำมานำมาคำนวณหาค่าเฉลี่ยของราคาประเมินที่ดิน
3. นำคอลัมน์ไร่ งาน ตารางวา มาคำนวณเปลี่ยนเป็นคอลัมน์ขนาดพื้นที่ต่อตารางวาโดย 1 ไร่ เท่ากับ 4 งาน และ 1 งาน เท่ากับ 100 ตารางวา
4. นำคอลัมน์ราคาประเมินและราคาขายมาหารด้วยขนาดพื้นที่ เปลี่ยนเป็นคอลัมน์ราคาประเมินต่อตารางวาและราคาขายต่อตารางวา

5.3 กระบวนการเตรียมข้อมูลและการสำรวจข้อมูลเบื้องต้น

1. กระบวนการเตรียมข้อมูล มุ่งเน้นการเตรียมข้อมูลให้อยู่ในสภาพพร้อมใช้ในการวิเคราะห์หรือการเรียนรู้เชิงข้อมูล โดยเปลี่ยนแปลงและปรับปรุงข้อมูลตามความจำเป็น เพื่อเพิ่มคุณภาพและความน่าเชื่อถือของข้อมูล และเพิ่มประสิทธิภาพในกระบวนการเรียนรู้ของเครื่อง
 - (1) ลบคอลัมน์ประเภททรัพย์สินที่เป็นห้องชุดออก
 - (2) ลบแถวที่มีค่าราคาขายได้/ราคาเสนอสูงสุดน้อยกว่า 0
 - (3) เปลี่ยนประเภทของคอลัมน์ ราคาประเมิน ราคาขายได้/ราคาเสนอสูงสุด ไร่ งาน ตารางวา

(4) นำคอลัมน์ประเภททรัพย์สินมาทำการเข้ารหัสแบบวัน-ฮอต (One-Hot Encoding) วิธีนี้ทำให้แต่ละค่าของข้อมูลประเภทกลายเป็นตัวแปรใหม่ที่มีค่า 0 หรือ 1 เพื่อระบุว่าข้อมูลอยู่ในหมวดหมู่ใดโดยในคอลัมน์ประเภททรัพย์สินจะมีอยู่ 2 ประเภทคือที่ดินพร้อมสิ่งปลูกสร้างและที่ดินว่างเปล่า

2. การสำรวจข้อมูลเบื้องต้น ดำเนินการศึกษาการกระจายตัวของราคาขายต่อตารางวา โดยพบว่าข้อมูลส่วนมากจะอยู่ในช่วง 1,812.95 บาทถึง 25,416.67 บาทและมีค่าข้อมูลผิดปกติ (Outliers) อยู่ 12 จุดในชุดข้อมูลที่มีค่ามากกว่า 60,355.03 บาทของขอบเขตสูงสุด (Upper Fence) จากนั้นทำการเปรียบเทียบจำนวนของประเภทที่ดิน โดยประเภททรัพย์สินที่มี 2 ประเภทคือที่ดินพร้อมสิ่งปลูกสร้างและที่ดินว่างเปล่า ข้อมูลทั้งหมด 193 รายการมีที่ดินพร้อมสิ่งปลูกสร้าง 127 รายการมากกว่าที่ดินว่างเปล่าที่มีเพียง 66 รายการ นอกจากนี้เมื่อศึกษาความสัมพันธ์ระหว่างขนาดพื้นที่ที่แสดงบนแกน X และราคาขายได้/ราคาเสนอสูงสุดที่แสดงบนแกน Y ดัง Figure 4.

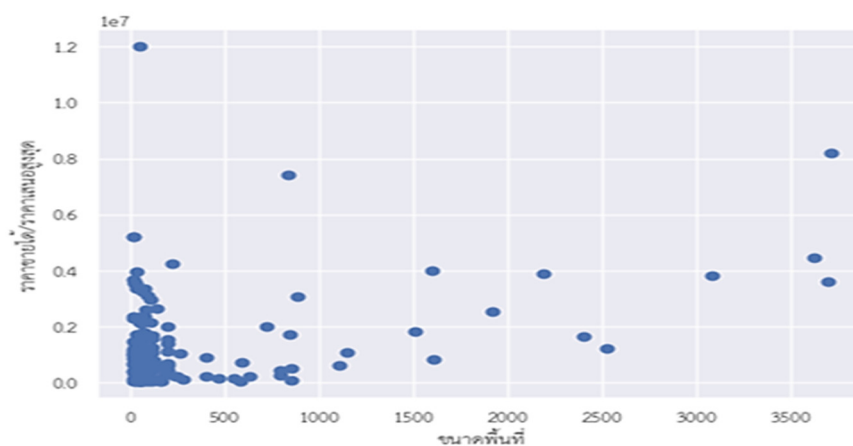


Figure 4. Relationship between area size and achievable selling price/ highest bid

จาก Figure 4. แสดงความสัมพันธ์ระหว่างขนาดพื้นที่ที่แสดงบนแกน X และราคาขายได้/ราคาเสนอสูงสุดที่แสดงบนแกน Y ในแผนภูมิจุดแบบกระจายนี้ข้อมูลกรณีศึกษาที่ดิน 193 แห่งที่ขายได้ส่วนมากจะมีขนาดพื้นที่ประมาณไม่เกิน 400 ตารางวา มีราคาขายได้/ราคาเสนอสูงสุดราคา 276 บาทถึงราคา 12,000,000 บาทและขนาดพื้นที่มากกว่า 400 ตารางวาถึงขนาดพื้นที่ไม่เกิน 4,000ตารางวาจะมีราคาขายได้/ราคาเสนอสูงสุดไม่เกิน ตารางวาจะมีราคาขายได้/ราคาเสนอสูงสุดไม่เกิน 9,000,000 บาท

5.4 การเรียนรู้ของเครื่อง

ในการวิจัยนี้ได้พัฒนาตัวแบบด้วยการใช้ภาษาไพทอนและไลบรารีหลักเช่น Scikit Learn เพื่อช่วยในการประเมินและพัฒนาประสิทธิภาพของตัวแบบ ข้อมูลถูกแบ่งออกเป็นสองส่วนโดย 80% ใช้สำหรับการฝึกสอนและ 20% เป็นข้อมูลทดสอบ สำหรับแบบจำลองที่ใช้ในการทดลอง ได้แก่ การถดถอยเชิงเส้น, ต้นไม้ถดถอย, แรนดอมฟอเรสต์ และเกรเดียนบูสทรี งานวิจัยนี้ได้ใช้วิธีการทดสอบแบบไขว้ทบ แบบ 5 โฟล (5-Fold Cross Validation) จากชุดข้อมูลฝึกสอน เพื่อค้นหาแบบจำลองที่มีประสิทธิภาพสูงสุด แบบจำลองเหล่านี้จะถูกปรับแต่งพารามิเตอร์โดยการใช้การค้นหาแบบกริด ซึ่งทำการทดลองกับทุกค่าพารามิเตอร์ที่กำหนดไว้ล่วงหน้า ดังแสดงใน Table 4. จากนั้นโมเดลที่ได้จากแต่ละอัลกอริทึมจะถูกนำไปทดสอบกับชุดข้อมูลทดสอบ 20% เพื่อประเมินประสิทธิภาพของตัวแบบในการทำนายข้อมูลจริง

Table 4. Parameters of the model using grid search.

Models	Parameters
Regression Trees	max_depth: [None, 10, 20, 30], min_samples_split: [2, 5, 10], min_samples_leaf: [1, 2, 4]
Random Forest	n_estimators: [100, 200, 300], max_depth: [None, 10, 20, 30], min_samples_split: [2, 5, 10], min_samples_leaf: [1, 2, 4]
Gradient boosted Trees	n_estimators: [100, 200, 300], learning_rate: [0.01, 0.1, 0.2] max_depth: [3, 4, 5], min_samples_split: [2, 5, 10] min_samples_leaf: [1, 2, 4]

ในการศึกษาปัจจัยที่ส่งผลต่อ ราคาขายต่อตารางวา งานวิจัยนี้ได้ทำการทดลองโดย ออกแบบการทดลองด้วยการจัดกลุ่มตัวแปรต้นเป็น 5 แบบการทดลอง ตามวงรอบของการปรับปรุงประสิทธิภาพเพื่อทำนายตัวแปรตาม ราคาขายต่อตารางวาโดยตัวแปรต้นในแต่ละการทดลองได้แสดงไว้ใน Table 5. ราคาประเมินต่อตารางวา ขนาดพื้นที่ ประเภททรัพย์สิน ราคาประเมินจากห้าตำแหน่ง และกลุ่มสถานที่สำคัญซึ่งจะมี 14 สถานที่โดยแบ่งเป็น สวนสาธารณะ โรงแรม ห้างสรรพสินค้า พิพิธภัณฑ์ สนามบิน โรงพยาบาล สถาบันการศึกษา ตลาด และศาสนสถานใช้หน่วยระยะทางเป็นกิโลเมตร

Table 5. The description of independent variable for each experimental.

Experiment	The appraisal price per square wah	Area size	Property type	Important location group	The appraisal price from 5 locations
1	✓	✓	✓	✓	
2	✓				
3	✓	✓	✓	✓	✓
4		✓	✓	✓	✓
5		✓	✓	✓	

5.5 การประเมินประสิทธิภาพ

เนื่องจากค่าข้อผิดพลาดที่ได้จะมีทั้งค่าบวกและค่าลบตรงจุดนี้เอง จึงทำให้ค่าความคลาดเคลื่อนเฉลี่ยสัมบูรณ์ ค่าความคลาดเคลื่อนเฉลี่ยรากที่สอง และค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวน เข้ามามีบทบาทสำคัญเนื่องจากทั้ง 3 วิธีจะมีการทำให้ค่าข้อผิดพลาดกลายเป็นค่าบวกก่อนเสมอ ก่อนที่จะนำค่าข้อผิดพลาดมารวมกันและหาค่าเฉลี่ยเพื่อประเมินประสิทธิภาพของโมเดลได้

6. ผลการวิจัย (Results)

จากวิธีการดำเนินงานข้างต้นสามารถสรุปผลการดำเนินงานในแต่ละการทดลองได้ดังนี้



Table 6. Results of the experiment.

Experiment/ Best Model	Hyper parameter	R^2	MAE	RMSE
Experiment 1 Gradient Boosted Trees	LR: 0.01 MD: 3 MSL: 2 MSS: 10 NE: 200	0.58	10911	22252
Experiment 2 Regression Tree	MD:None MSL: 2 MSS: 2	0.73	8908	17822
Experiment 3 Regression Tree	MD: 10 MSL: 1 MSS: 2	0.65	10482	20207
Experiment 4 Gradient Boosted Trees	LR: 0.2 MD: 3 MSL: 1 MSS: 10 NE: 100	0.80	7929	15281
Experiment 5 Gradient Boosted Trees	LR: 0.01 MD: 3 MSL: 1 MSS: 10 NE: 300	0.74	10807	17400

Table 6. แสดงผลลัพธ์ที่ได้จากการทำโมเดลซีเลคชั่น ของการค้นหาค่าของไฮเปอร์พารามิเตอร์ของโมเดลแบบกริด ซึ่งจาก Table 6. สามารถอธิบายตัวย่อของตัวแปรต่าง ๆ ได้ดังนี้ (Pedregosa et al., 2012) พารามิเตอร์ (Hyperparameter) มี LR, MD, MSL, MSS, NE คือ อัตราการเรียนรู้ (Learning Rate, LR) กำหนดค่าที่ใช้ในกระบวนการฝึกสอนของเกรเดียนบูสทรี ค่าที่เลือกจะมีผลต่อการ อัปเดตค่าของแต่ละต้นไม้ในเอนเซมเบิล (Ensemble) ในแต่ละรอบ ค่าความลึกของต้นไม้ (Max Depth, MD) กำหนดความลึกสูงสุดของต้นไม้ค่ามากมักทำให้ต้นไม้ซับซ้อนมากขึ้น แต่หากมีค่ามากเกินไปอาจเป็นการโอเวอร์ฟิตติ้ง มินแซมเพิลลีฟ (Min Samples Leaf, MSL) กำหนดจำนวนตัวอย่างขั้นต่ำที่ต้องอยู่ในใบ (Leaf Node) ค่านี้ช่วยลดความซับซ้อนของต้นไม้และป้องกันการโอเวอร์ฟิตติ้ง มินแซมเพิลสปลิท (Min Samples Split, MSS) จำนวนตัวอย่างขั้นต่ำในการแยกและ เอ็นเอสติเมท (N Estimators, NE) จำนวนต้นไม้ ตามลำดับ

การทดลองที่ 1 ได้ผลการคัดเลือกโมเดลโดยพบว่า เกรเดียนบูสทรีเป็นโมเดลที่มีประสิทธิภาพดีที่สุด และจากการวิเคราะห์ความสำคัญของคุณลักษณะโดยใช้ค่าสารสนเทศ (Information Gain) แสดงว่าคุณลักษณะ “ราคาประเมินต่อตารางวา” มีผลต่อการทำนายมากที่สุด รองลงมาจะเป็นกลุ่มสถานที่สำคัญขนาดพื้นที่ และประเภททรัพย์สินที่ไม่มีผลต่อการทำนาย

การทดลองที่ 2 ได้สังเกตเห็นถึงความสำคัญของคุณลักษณะ “ราคาประเมินต่อตารางวา” ที่มีผลต่อการทำนายมากที่สุดจากการวิเคราะห์ผลการทดลองที่ 1 ดังนั้นจึงนำเฉพาะส่วนราคาประเมินต่อตารางวามาทดสอบเพิ่มเติม โดยผลการคัดเลือกโมเดลพบว่า โมเดลต้นไม้ตัดสินใจแบบถดถอย เป็นโมเดลที่มีประสิทธิภาพดีที่สุด ซึ่งค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวนของโมเดลที่มีค่าเพิ่มขึ้น ค่าความคลาดเคลื่อนเฉลี่ยสมบูรณ์ และค่าความคลาดเคลื่อนเฉลี่ยรากที่สองมีค่าน้อยลงเมื่อเทียบกับการทดลองที่ 1

การทดลองที่ 3 ได้เพิ่มตัวแปรราคาประเมินจากพื้นที่ใกล้เคียงห้าตำแหน่งเข้าไปในการทดลองที่ 1 เพื่อตรวจสอบว่าราคาประเมินจากห้าตำแหน่งมีผลต่อการทำนายมากขึ้นหรือไม่ ซึ่งผลการคัดเลือกโมเดลพบว่า ต้นไม้ตัดสินใจแบบถดถอย เป็นโมเดลที่มีประสิทธิภาพดีที่สุด และจากการวิเคราะห์ความสำคัญของคุณลักษณะ พบว่าโมเดลให้เห็นว่า “ราคาประเมินต่อตารางวา” ยังคงมีผลต่อการทำนายมากที่สุด รองลงมาคือกลุ่มสถานที่สำคัญ และขนาดพื้นที่ ในขณะที่ราคาประเมินจากพื้นที่ใกล้เคียงห้าตำแหน่งและประเภททรัพย์สินไม่มีผลต่อการทำนาย เมื่อเทียบกับการทดลองที่ 1 ผลลัพธ์ของการทดลองที่ 3 มีประสิทธิภาพดีกว่า แต่ยังคงด้อยกว่าโมเดลที่ได้จากการทดลองที่ 2

การทดลองครั้งที่ 4 และ 5 เป็นการนำราคาประเมินต่อตารางวาออกจากการทดลองที่ 3 และ 1 ตามลำดับเพื่อทดสอบว่าถ้าไม่ทราบราคาประเมินต่อตารางวาของที่ดินผืนนั้น ๆ จะสามารถทำการประเมินราคาที่ใกล้เคียงราคาขายจริงได้อยู่หรือไม่

การทดลองที่ 4 ได้โมเดลเกรเดียนบูสทรีเป็นโมเดลที่มีประสิทธิภาพดีที่สุด และจากการวิเคราะห์ความสำคัญของคุณลักษณะ พบว่าโมเดลให้ความสำคัญกับคุณลักษณะกลุ่มสถานที่สำคัญมีผลต่อการทำนายมากที่สุด รองลงมาคือขนาดพื้นที่ ราคาประเมินจากห้าตำแหน่ง และประเภทที่ดิน นอกจากนี้สามารถศึกษาค่าประสิทธิภาพของโมเดลต่าง ๆ ในการคัดเลือกโมเดลที่ดีที่สุด ผลการวิจัย ดัง Table 7.

Table 7. Performance metrics of various models for selecting the best model from Experiment 4.

โมเดล	R^2	MAE	RMSE
Linear Regression	0.23	17378.70	30199.13
Regression Tree	0.21	13539.32	30596.87
Random Forest	0.72	10371.91	18180.34
GradientBoosting	0.74	10807.34	17400.46

จาก Table 7. พบว่า ค่าประสิทธิภาพของโมเดลต่าง ๆ ที่ได้จากการทำการค้นหาพารามิเตอร์แบบกริดของการทดลองที่ 4 โดยจากตารางพบว่า ค่าสัมประสิทธิ์การอธิบายความแปรปรวน (R^2) ของเกรเดียนบูสทรีสูงสุดที่ 0.7445 ซึ่งแสดงว่าโมเดลสามารถอธิบายความแปรปรวนในข้อมูลได้ประมาณ 74% ค่าความคลาดเคลื่อนเฉลี่ยสมบูรณ์ ของเกรเดียนบูสทรีอยู่ที่ 10807.34 ซึ่งสูงกว่าค่าของแรนดอมฟอเรสต์แต่ยังคงค่อนข้างต่ำ นอกจากนี้ความคลาดเคลื่อนเฉลี่ยรากที่สอง เกรเดียนบูสทรีอยู่ที่ 17400.46 ซึ่งต่ำที่สุดในบรรดาโมเดลทั้งหมด แสดงว่าโมเดลมีข้อผิดพลาดในการทำนายต่ำที่สุด

การทดลองที่ 5 พบว่าเกรเดียนบูสทรีเป็นโมเดลที่มีประสิทธิภาพดีที่สุด เมื่อทำการตัดคุณลักษณะด้านราคาประเมินออกจากตัวแปรต้น โมเดลที่ได้มีประสิทธิภาพเป็นอันดับ 2 รองจากโมเดลที่ได้จากการทดลองที่ 4 จากการวิเคราะห์ความสำคัญของคุณลักษณะพบว่าโมเดลให้ความสำคัญกับกลุ่มสถานที่สำคัญมากที่สุด รองลงมาคือขนาดพื้นที่ และประเภทที่ดิน แสดงให้เห็นว่าสามารถพยากรณ์ราคาขายต่อตารางวาได้ใกล้เคียงกับราคาขายจริงหากทราบว่าที่ดินที่ต้องการประเมินนั้นอยู่ห่างจากสถานที่สำคัญเท่าใดเมื่อเทียบกับ การประเมินโดยทราบว่าขนาดของพื้นที่ หรือประเภททรัพย์สินเป็นอะไร

ข้อสังเกตที่น่าสนใจจากการเปรียบเทียบผลการทดลองที่ 3 และผลการทดลองที่ 4 ที่ตัดตัวแปรราคาประเมินต่อตารางวาของที่ดินผืนนั้น ๆ ออกพบว่าได้โมเดลที่มีประสิทธิภาพสูงขึ้น ซึ่งอาจจะเป็นผลมาจากการโอเวอร์ฟิตของโมเดลที่ให้น้ำหนักไปในการจดจำราคาประเมินต่อตารางวาเป็นหลัก หรือ เมื่อฟีเจอร์ที่มีอิทธิพลถูกลบออก โมเดลจะพึ่งพาฟีเจอร์อื่น ๆ ที่อาจมีข้อมูลที่สำคัญมากขึ้น ซึ่งทำให้โมเดลสามารถจับโครงสร้างที่แท้จริงของข้อมูลได้อย่างมีประสิทธิภาพมากขึ้น

การนำโมเดลมาทดลองใช้ในเว็บไซต์

เมื่อได้โมเดลเป็นที่เรียบร้อยแล้ว ผู้วิจัยได้ทำการออกแบบและพัฒนาเว็บไซต์สำหรับการทำนายราคาขายต่อตารางวาโดยใช้โมเดลเกรเดียนบูสท์ของการทดลองที่ 4 ที่มีค่าผลการทดลองที่ดีที่สุดในการทำนายข้อมูล เมื่อทดลองกรอกข้อมูลค่าละติจูด ลองติจูด ไร่ งาน ตารางวา และเลือกประเภทที่ดินว่างเปล่าโดยที่ดินที่เลือกมาจากระบบค้นหารูปแปลงที่ดินของกรมที่ดิน ที่เป็นจังหวัดขอนแก่น อำเภอเมืองขอนแก่น ตำบลในเมือง พบผลการวิจัย ดัง Figure 5.

Figure 5. Entering data to predict the price per square meter.

จาก Figure 5. พบว่า ในเนื้อที่ 0 ไร่ 0 งาน 49.8 ตารางวา ราคาประเมินที่ดินกรมธนารักษ์ 6,500 บาทต่อตารางวา ค่าพิกัดแปลง 16.44714910,102.82108872 ผลลัพธ์ที่ได้จากทำนายราคาขายต่อตารางวาของแปลงที่ดินนี้คือ 7290.63 บาทต่อตารางวา

7. สรุปผลการวิจัย (Conclusion)

การวิจัยนี้ได้ทำการพัฒนาอัลกอริทึมการเรียนรู้ของเครื่องเพื่อประเมินราคาที่ดินในเขตอำเภอเมือง จังหวัดขอนแก่น โดยใช้ข้อมูลจากกรมบังคับคดีและกรมที่ดิน จากการทดลองปรับชุดกลุ่มตัวแปรต้นและการทำการคัดเลือกตัวแปรด้วยกระบวนการค้นหาพหุมิเตอร์แบบกริด และการทำการทดสอบแบบไขว้ทบ ผลการวิจัยพบว่าโมเดลเกรเดียนบูสท์ให้ผลประสิทธิภาพการทำนายราคาที่ดินจากข้อมูลการซื้อขายจากข้อมูลกรมบังคับคดีได้ดีที่สุด โดยมีค่าสัมประสิทธิ์กำลังสองของการอธิบายความแปรปรวนสูงที่สุดมีค่า 0.80 ค่าความคลาดเคลื่อนเฉลี่ยสมบูรณ์ และค่าความคลาดเคลื่อนเฉลี่ยรากที่สองมีน้อยที่สุดโดยมีค่า 7929.40 และ 15281.33 ตามลำดับ ซึ่งจากโมเดลที่ได้ อนุมาณความสำคัญของคุณลักษณะ จากมากไปน้อยได้ดังต่อไปนี้ กลุ่มสถานที่สำคัญขนาดพื้นที่ ราคาประเมินจากห้าตำแหน่ง และประเภททรัพย์สิน ข้อสังเกตที่น่าสนใจอีกประการคือ ถึงแม้ว่าราคาประเมินของ ที่ดินนั้น ๆ จะมีความสำคัญของคุณลักษณะที่สูงที่สุดจากทุก ๆ ตัวแปร แต่เมื่อตัดออกทำให้ได้โมเดลที่มีประสิทธิภาพสูงขึ้น ซึ่งอาจเกิดจากการโอเวอร์ฟิตของตัวแบบที่อาศัยการจดจำราคาประเมิน หรือ

การตัดตัวแปรที่มีอิทธิพลส่งผลให้โมเดลที่พาดูคุณลักษณะอื่น ๆ ที่มีข้อมูลสำคัญมากขึ้น ทำให้โมเดลสามารถระบุและหาความสัมพันธ์ที่แท้จริงของข้อมูลได้อย่างมีประสิทธิภาพมากยิ่งขึ้น

8. ข้อเสนอแนะงานวิจัย (Recommendation)

สำหรับการพัฒนาต่อยอดงานวิจัยนี้ สามารถที่จะพัฒนาในส่วนของการเก็บฐานข้อมูลเพิ่มเติมในส่วนของการติดตามของซึ่งยังมีข้อจำกัดคือ ข้อมูลจากกรมบังคับคดีจะถูกปล่อยออกไปเมื่อที่ดินแปลงนั้นขายได้ทำให้ไม่สามารถเก็บข้อมูลในส่วนนี้ย้อนหลังได้ และสามารถเพิ่มคุณลักษณะตัวแปรต้นในการนำข้อมูลรูปถ่ายของเมืองตามระยะเวลามาทำการจำแนกรูปภาพเชิงความหมาย (Semantic Segmentation) เพื่อสกัดคุณลักษณะของพื้นที่ ที่มีการเปลี่ยนแปลงเชิงเวลาเพิ่มเติมสำหรับโมเดลในการพยากรณ์เป็นแนวทางในการพัฒนาต่อไป

9. กิตติกรรมประกาศ (Acknowledgement)

งานวิจัยชิ้นนี้ได้รับการสนับสนุนทรัพยากร คอมพิวเตอร์จากบริษัทอินเทอร์เน็ตประเทศไทย ซึ่งเป็นส่วนสำคัญในการช่วยเหลือและสนับสนุนสำหรับการวิจัย

10. เอกสารอ้างอิง (References)

- Bergstra, J., & Bengio, Y. (2012). Random Search for Hyper-Parameter Optimization. *Journal of Machine Learning Research*, 13(10), 281–305.
- Burnham, K. P., & Anderson, D. R. (2002). *Model Selection and Multimodel Inference: a Practical Information-Theoretic Approach*. Springer Science & Business Media.
- Chandrashekar, G., & Sahin, F. (2014). A Survey on Feature Selection Methods. *Computers & Electrical Engineering*, 40(1), 16–28. <https://doi.org/10.1016/j.compeleceng.2013.11.024>.
- Choompol, A. (2019). *Feature Selection and Redundant Feature Elimination for Opinion Classification on Social Network*. [Doctoral dissertation, Mahasarakham University]. Mahasarakham University Intellectual Repository. <http://202.28.34.124/dspace/handle/123456789/537>. (In Thai).
- Guyon, I., & Elisseeff, A. (2003). An Introduction to Variable and Feature Selection. *Journal of Machine Learning Research*, 3(7–8), 1157–1182.
- Kohavi, R. (1995). A Study of Cross-Validation and Bootstrap for Accuracy Estimation and Model Selection. *Proceedings of the 14th International Joint Conference on Artificial Intelligence - Volume 2*, 1137–1143.
- Kumpu, P., & Piyathamronchai, K. (2024). Application of the Geographic Information System to Develop Land Valuation Models, Case Study: Muang Chiangmai District, Chiangmai Province. *The Journal of Spatial Innovation Development*, 5(2), 42–64. (In Thai)
- Lundberg, S., & Lee, S.-I. (2017). A Unified Approach to Interpreting Model Predictions. *31st Conference on Neural Information Processing Systems*. 4765–4774. <https://doi.org/10.48550/ARXIV.1705.07874>.

- Na Bangchang, K. (2011). *A Variable Selection in Multiple Linear Regression Models Based on Tabu Search* [National Institute of Development Administration].
<https://doi.org/10.14457/NIDA.the.2011.17>. (In Thai).
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Müller, A., Nothman, J., Louppe, G., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., & Duchesnay, É. (2012). Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research*, 12, 2825–2830.
<https://doi.org/10.48550/ARXIV.1201.0490>
- Soltani, A., Heydari, M., Aghaei, F., & Pettit, C. J. (2022). Housing Price Prediction Incorporating Spatio-Temporal Dependency into Machine Learning Algorithms. *Cities*, 131, 103941.
<https://doi.org/10.1016/j.cities.2022.103941>.
- Vanderplas, J. (2017). *Python Data Science Handbook: Essential Tools for Working with Data*. O'Reilly.

Development of Mathematics Achievement, Mathematics Problem- Solving Ability, and Responsibility Using Inquiry-Based Learning Management of Grade, 7 Students การพัฒนาผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ ความสามารถในการแก้ปัญหาทางคณิตศาสตร์ และ ความรับผิดชอบโดยใช้การจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1

Chayanit Muangdoo^{1,*}, and Yupadee Panarach²

ชยานิต เมืองดู่^{1,*}, และ ยุภาดี ปณาราช²

Received: 3 March 2024;

Revised: 13 June 2024;

Accepted: 16 June 2024;

Published: 15 August 2024;

Abstract

This study aimed to 1) compare the mathematics achievement before and after the inquiry-based learning management, 2) compare the mathematics achievement after learning management with the 70 percent criterion, and 3) develop the level of mathematics problem-solving and responsibility after the learning management of grade, 7 students. The instruments used for data collection were learning management plans, achievement tests, problem-solving ability tests, and responsibility tests. Data analysis was analyzed using mean, standard deviation, the t-test. The results found that 1) the students have higher mathematics achievement after learning than before learning, and higher than the 70 percent criterion at a statistically significant of 0.05 levels, 2) problem-solving ability at high levels, and 3) responsibility after learning management at most levels.

Keywords: Mathematics Achievement, Mathematics Problem-Solving Ability, Responsibility, Inquiry-Based Learning Management.

บทคัดย่อ

การศึกษานี้มีวัตถุประสงค์เพื่อ 1) เปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ ก่อนและหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ 2) เปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์หลังการจัดการเรียนรู้กับเกณฑ์ร้อยละ 70 3) พัฒนาระดับการแก้ปัญหาทางคณิตศาสตร์และความรับผิดชอบหลังการจัดการเรียนรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1 เครื่องมือที่ใช้เก็บข้อมูล ได้แก่ แผนการจัดการเรียนรู้

^{1,*} Student, Faculty of Education, Udon Thani Rajabhat University, Udon Thani 41000, Thailand; นักศึกษา คณะครุศาสตร์ มหาวิทยาลัยราชภัฏอุดรธานี จังหวัดอุดรธานี 41000 ประเทศไทย; Email: pratumtetc@gmail.com

² Associate Professor, Dr., Faculty of Education, Udon Thani Rajabhat University, Udon Thani 41000, Thailand; รองศาสตราจารย์ ดร. คณะครุศาสตร์ มหาวิทยาลัยราชภัฏอุดรธานี จังหวัดอุดรธานี 41000 ประเทศไทย; Email: yupadee.p@udru.ac.th

*Corresponding authors: Chayanit Muangdoo (pratumtetc@gmail.com)



แบบทดสอบวัดผลสัมฤทธิ์ แบบวัดความสามารถในการแก้ปัญหา แบบวัดความ
รับผิดชอบ การวิเคราะห์ข้อมูลโดยใช้ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐาน การ
ทดสอบค่าที ผลการวิจัยพบว่า 1) นักเรียนมีผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์หลัง
เรียนสูงกว่าก่อนเรียน และสูงกว่าเกณฑ์ร้อยละ 70 อย่างมีนัยสำคัญทางสถิติที่ระดับ
0.05 2) ความสามารถในการแก้ปัญหายู่ในระดับมาก และ 3) ความรับผิดชอบหลัง
การจัดการเรียนรู้ในระดับมากที่สุด

คำสำคัญ: ผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์, ความสามารถในการแก้ปัญหาด้าน
คณิตศาสตร์, ความรับผิดชอบ, การจัดการเรียนรู้แบบสืบเสาะหาความรู้

1. บทนำ (Introduction)

ปัจจุบันประเทศไทยอยู่ในยุคไทยแลนด์ 4.0 มีการขับเคลื่อนประเทศด้วยเทคโนโลยีนวัตกรรม และความคิด
สร้างสรรค์ เพื่อนำประเทศไทยไปสู่ความมั่นคง มั่งคั่ง และยั่งยืนได้ (Chularut, 2018) ในส่วนของภาคการจัดการศึกษา
จำเป็นอย่างยิ่งที่ต้องมีการวางแผน และการปฏิรูปกระบวนการจัดการเรียนการสอนให้เข้ากับนักเรียน เพื่อพัฒนาให้ตอบสนองต่อ
ความต้องการของผู้เรียน เนื่องจากการศึกษาในยุคนี้ ไม่ใช่เป็นเพียงการให้ความรู้ระหว่างครูผู้สอนกับนักเรียนเท่านั้น
นอกจากความรู้ที่ผู้เรียนจะได้รับแล้ว ผู้เรียนยังจะต้องได้รับการพัฒนาทักษะที่สำคัญในการดำรงชีวิตด้วย (Chularut, 2018)
การจัดการเรียนรู้ในยุคนี้เป็นการจัดการเรียนรู้ที่มุ่งสอนให้ผู้เรียนมีคุณภาพชีวิตที่ดีขึ้นและเป็นมนุษย์ที่สมบูรณ์และให้
ตอบสนองต่อความต้องการของสังคมได้ ซึ่งเมื่อศึกษาข้อมูลในเรื่องดังกล่าวจะพบว่าการจัดการเรียนรู้ในยุคนี้ สามารถ
ตอบสนองต่อการพัฒนาผู้เรียนให้เป็นไปตามกระแสสังคมโลกและศตวรรษที่ 21 ได้เป็นอย่างดี (Jitchayawanit, 2019)
ดังนั้น การปรับเปลี่ยนวิธีการสอนจึงเป็นประเด็นสำคัญในการพัฒนาผู้เรียนให้มีประสิทธิภาพสามารถส่งเสริมให้ผู้เรียน
สามารถนำความรู้ไปประยุกต์ในชีวิตประจำวันและในการศึกษาต่อได้อย่างมีประสิทธิภาพ (Research Administration
and Educational Quality Assurance Division, 2017)

กระบวนการจัดการเรียนการสอนแบบสืบเสาะหาความรู้จึงนับเป็นการจัดการเรียนรู้หนึ่ง ที่เน้นให้ผู้เรียนสามารถ
พัฒนาความสามารถด้านการแก้ปัญหา และการค้นคว้าหาความรู้โดยจะให้ผู้สอนตั้งคำถามและให้ผู้เรียนศึกษาค้นคว้าหา
คำตอบด้วยตนเอง (Sutthiratt, 2015) เนื่องจากการจัดการเรียนรู้แบบสืบเสาะหาความรู้ เป็นการจัดการเรียนรู้ที่ส่งเสริมให้
ผู้เรียนแสวงหาความรู้ด้วยตนเอง โดยใช้วิธีการและทักษะกระบวนการทางวิทยาศาสตร์ ผ่านการที่ผู้สอนกระตุ้นให้ผู้เรียน
เกิดคำถาม เกิดความคิด และลงมือแสวงหาความรู้ เพื่อให้ผู้เรียนสามารถศึกษาค้นคว้าหาความรู้ได้ด้วยตนเองโดยที่
ผู้สอนจะเป็นผู้ชี้แนะหรือให้คำแนะนำ (Khammani, 2010) การสืบเสาะหาความรู้จึงเป็นการค้นหาคำตอบที่สนใจ ผ่านการ
ทำงานอย่างเป็นระบบ รอบคอบ แต่มีอิสระ การสืบเสาะหาความรู้จะมีลำดับขั้นตอนที่ไม่ตายตัวและจะเป็นการศึกษาหา
คำตอบเดิม ๆ ซ้ำ ๆ อยู่หลายรอบ ซึ่งอาจเกิดคำถามขึ้นมาใหม่ที่ต้องสืบเสาะหาคำตอบต่อไป หมุนวนเป็นวัฏจักร 5 ขั้นตอน
ประกอบไปด้วย ขั้นที่ 1 การสร้างความสนใจ ขั้นที่ 2 การสำรวจและค้นหา ขั้นที่ 3 การอธิบายและลงข้อสรุป ขั้นที่ 4 การ
ขยายความรู้ และขั้นที่ 5 การประเมินความรู้ (Institute for the Promotion of Teaching Science and Technology,
2019) แต่รูปแบบการจัดการเรียนรู้แบบสืบเสาะหาความรู้เพียงอย่างเดียว ยังไม่เพียงพอให้เป็นไปตามจุดมุ่งหมายของการ
จัดการศึกษาได้ เนื่องจากการจัดการศึกษาในยุคไทยแลนด์ 4.0 ต้องมีการนำเทคโนโลยีใหม่ ๆ เข้ามาประยุกต์ใช้ในการ
ปฏิบัติงาน เพื่อปรับเปลี่ยนให้เหมาะสมและสอดคล้องกับยุคสมัย และให้ผู้เรียนมีความสามารถในการศึกษาค้นคว้าหาความรู้
ให้เกิดผลได้อย่างมีประสิทธิภาพ

ดังนั้น ผู้วิจัยจึงสนใจที่จะนำการจัดการเรียนรู้แบบสืบเสาะหาความรู้ มาใช้จัดกิจกรรมการเรียนรู้ เรื่องสมการเชิงเส้นสองตัวแปร กับนักเรียนชั้นมัธยมศึกษาปีที่ 1 ซึ่งวิธีการจัดการเรียนรู้ดังกล่าว จะเป็นแนวทางเพื่อแก้ไขปัญหาคุณภาพการจัดการศึกษา พัฒนารูปแบบการเรียนรู้ พัฒนาผลสัมฤทธิ์ทางการเรียน ความสามารถในการแก้ปัญหาทางคณิตศาสตร์และความรับผิดชอบของผู้เรียน ผู้วิจัยจึงต้องการศึกษาว่าการจัดการเรียนรู้แบบสืบเสาะหาความรู้ จะทำให้ผู้เรียนมีผลสัมฤทธิ์ทางการเรียนหลังสูงกว่าก่อนเรียนหรือไม่ ความสามารถในการแก้ปัญหาทางคณิตศาสตร์และความรับผิดชอบจะเป็นอย่างไร เพื่อนำผลการวิจัยที่ได้มาเป็นแนวทางในการจัดการเรียนรู้วิชาคณิตศาสตร์ของโรงเรียนต่อไป

2. วัตถุประสงค์งานวิจัย (Research Objectives)

1. เพื่อเปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ ก่อนและหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1
2. เพื่อพัฒนาระดับความสามารถในการแก้ปัญหาทางคณิตศาสตร์หลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1
3. เพื่อศึกษาความรับผิดชอบหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1

3. สมมติฐานงานวิจัย (Research Hypothesis)

1. นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ หลังสูงกว่าก่อนการจัดการเรียนรู้แบบสืบเสาะหาความรู้
2. นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีความสามารถในการแก้ปัญหาทางคณิตศาสตร์หลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ในระดับมาก
3. นักเรียนมีความรับผิดชอบหลังการจัดการเรียนรู้โดยรวมอยู่ในระดับมากที่สุด

4. กรอบแนวคิดงานวิจัย (Conceptual Framework)

งานวิจัยนี้กำหนดกรอบแนวคิดในการทำวิจัยและดำเนินการตามแนวคิดทฤษฎี ดัง Figure 1.

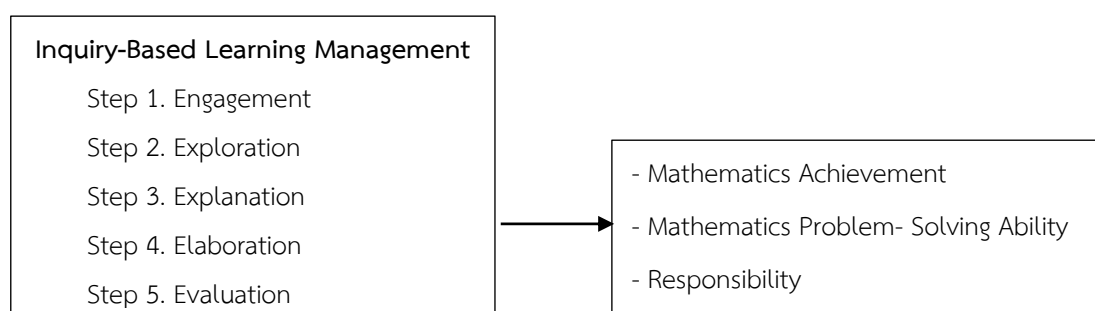


Figure 1. Conceptual Framework

5. การทบทวนวรรณกรรมและทฤษฎีที่เกี่ยวข้อง (Literature Review)

5.1 ความหมายของการจัดการเรียนรู้แบบสืบเสาะหาความรู้

Sinthapanon et al. (2002) กล่าวสรุปได้ว่า นักการศึกษาทั้งต่างประเทศและในประเทศไทยได้แก่ Sucman, Gagne, Sun และ Trowbridge, อีริช บรูณโซติ และวีระยุทธ วิเชียรโชติ ได้ให้ความหมายของวิธีการสอนแบบสืบเสาะหาความรู้ว่าเป็นวิธี ที่พัฒนาความสามารถในการคิด การแก้ปัญหาหรือแสวงหาความรู้โดยใช้กระบวนการทางความคิดเพื่อแสวงหาความรู้และค้นพบคำตอบด้วยตนเอง

Khammani (2010) และ Chaiyo & Sakolkeart (2019) ได้ให้คำนิยามของการจัดการเรียนการสอนโดยเน้นกระบวนการสืบเสาะหาความรู้ไว้พอสรุปได้ว่า หมายถึง กระบวนการที่ผู้สอนจัดสถานการณ์กระตุ้นให้นักเรียนเกิดความคิดและดำเนินการค้นคว้าหาความรู้ด้วยตนเอง จากนั้นนำผู้เรียนสู่การสรุปผลหรือการหาคำตอบด้วยตนเอง ผู้สอนมีหน้าที่ให้ความช่วยเหลือและสนับสนุนผู้เรียนตลอดทั้งกระบวนการจัดการเรียนรู้

5.2 ประเภทของการจัดการเรียนรู้แบบสืบเสาะหาความรู้

1. ผู้สอนมีบทบาทสำคัญในการสืบเสาะหาความรู้ วิธีนี้ผู้สอนมีบทบาทสำคัญในการใช้คำถามกระตุ้นให้เป็นแนวทางให้ผู้เรียนคิดหาคำตอบ เหมาะสำหรับการเริ่มสอนแบบสืบเสาะหาความรู้ เนื่องจากผู้สอนจะเป็นผู้ใช้คำถามนำไปสู่คำตอบและพยายามกระตุ้นให้ผู้เรียนตั้งคำถามอยู่เสมอ โดยผู้สอนจะเป็นผู้ตั้งคำถามเป็นส่วนใหญ่

2. ผู้สอนและผู้เรียนร่วมกันในการสืบเสาะหาความรู้ วิธีนี้ผู้สอนและผู้เรียนเป็นผู้ดำเนินการในการสืบเสาะหาความรู้ร่วมกัน โดยผู้สอนเป็นผู้ตั้งคำถามเท่าๆ

3. ผู้เรียนเป็นผู้มีบทบาทสำคัญในการสืบเสาะหาความรู้ การสอนแบบนี้นักเรียนจะต้องเป็นผู้ตั้งคำถามและตอบคำถามเป็นส่วนใหญ่ หลังจากที่ได้ฝึกการตั้งคำถามและตอบคำถามจนคุ้นเคยมาแล้ว ผู้เรียนได้รับการพัฒนาความคิด การตั้งคำถามในกระบวนการสืบเสาะเพื่อหาคำตอบด้วยตามลำดับขั้น

5.3 ขั้นตอนการจัดการเรียนรู้แบบสืบเสาะหาความรู้

Sinthapanon et al. (2002) สรุปขั้นตอนการจัดกิจกรรมการเรียนการสอนแบบสืบเสาะหาความรู้โดยแบ่งเป็น 5 ขั้นตอน ดังนี้

1. สอนสร้างสถานการณ์หรือปัญหาจากเนื้อหาหลักสูตรให้สอดคล้องกับจุดประสงค์การเรียนรู้เป็นการนำเข้าสู่บทเรียนด้วยปัญหาเพื่อกระตุ้นให้ผู้เรียนคิดและแก้ปัญหา การนำเข้าสู่บทเรียนอาจทำได้หลายวิธี

2. ใช้คำถามในการอธิบายเพื่อนำไปสู่แนวทางในการหาคำตอบ การใช้คำถามนี้จะต้องอาศัยสถานการณ์หรือปัญหาที่กำหนดขึ้นโดยใช้คำถามเป็นชุดต่อเนื่องสัมพันธ์กัน

3. ใช้คำถามเพื่อนำไปสู่การออกแบบกำหนดวิธีการศึกษา การทดลองเพื่อหาคำตอบคำถามในขั้นนี้เป็นคำถามเพื่อนำไปสู่การอธิบายวิธีการหาความรู้

4. ดำเนินการศึกษาค้นคว้าสืบเสาะ หาความรู้ ผู้สอนต้องใช้คำถามกระตุ้นให้ผู้เรียนได้ทำความเข้าใจในขั้นตอนการปฏิบัติกิจกรรมตามวิธีการที่ได้เลือกไว้ให้ชัดเจน จัดบันทึกข้อมูลไว้

5. ขั้่นอธิบายเพื่อสรุปผล ในขั้นนี้เป็นการใช้คำถามโดยอาศัยข้อมูลที่ได้จากการศึกษากันคว้าและการตอบคำถามเป็นหลักเพื่อนำไปสู่การสรุปหาคำตอบ

Institute for the Promotion of Teaching Science and Technology (2019) ได้เสนอขั้นตอนการสอนแบบสืบเสาะหาความรู้ (Inquiry Cycle) หรือแบบ 5E ประกอบด้วย 5 ขั้นตอน ดังนี้

1. ขั้นสร้างความสนใจ (Engagement) ในขั้นตอนนี้ผู้สอนจะนำผู้เรียนเข้าสู่บทเรียนหรือหัวข้อที่ต้องการจะสอน ซึ่งหัวข้อหรือเนื้อหาดังกล่าวอาจเกิดจากความสงสัยหรืออาจเริ่มจากตัวผู้เรียน ผู้สอนอาจกำหนดจากสถานการณ์ที่กำลังเกิดขึ้นในขณะนั้น ผู้สอนจะต้องตัวกระตุ้นให้ผู้เรียนเกิดความต้องการที่จะตั้งคำถามหรือประเด็นที่ต้องการจะศึกษา
2. ขั้นสำรวจและค้นหา (Exploration) หลังจากที่ได้เข้าใจประเด็นคำถามแล้ว ต่อมาจะต้องมีการกำหนดแนวทางในการตรวจสอบหรือการกำหนดสมมติฐานขึ้น เพื่อให้เกิดแนวทางที่เป็นไปได้ในการแก้ไขปัญหา หรือหาคำตอบ ผู้เรียนจะต้องเป็นผู้รวบรวมข้อมูลที่เกิดขึ้น
3. ขั้นอธิบายและลงข้อสรุป (Explanation) นำข้อมูลที่ได้เก็บรวบรวมข้อมูลมาวิเคราะห์ แปลผล สรุปผล และนำเสนอผลในรูปแบบต่าง ๆ
4. ขั้นขยายความรู้ (Elaboration) ผู้เรียนจะต้องสามารถนำความรู้ที่สร้างขึ้นมาเชื่อมโยงกับความรู้เดิม หรือขยายเพิ่มเติม ทั้งนี้เพื่อให้สามารถอธิบายสิ่งที่เกิดขึ้นจากผลการศึกษาให้เข้าใจมากขึ้น
5. ขั้นประเมิน (Evaluation) ดำเนินการประเมินการเรียนรู้ของผู้เรียนเพื่อตรวจสอบความรู้ ความเข้าใจ หลังจากผ่านกระบวนการเรียนรู้ จากนั้นควรมีการแนะนำแนวทางสู่การประยุกต์ใช้ความรู้ไปยังเรื่องอื่น ๆ

สถาบันส่งเสริมการสอนวิทยาศาสตร์และเทคโนโลยี นำแนวคิดทฤษฎีนี้ออกเผยแพร่แก่ครูโดยการจัดการอบรมการเรียนรู้การสอนแบบสืบเสาะหาความรู้ทั่วประเทศ สามารถอธิบายภาพดัง Figure 2. (Institute for the Promotion of Teaching Science and Technology, 2019)

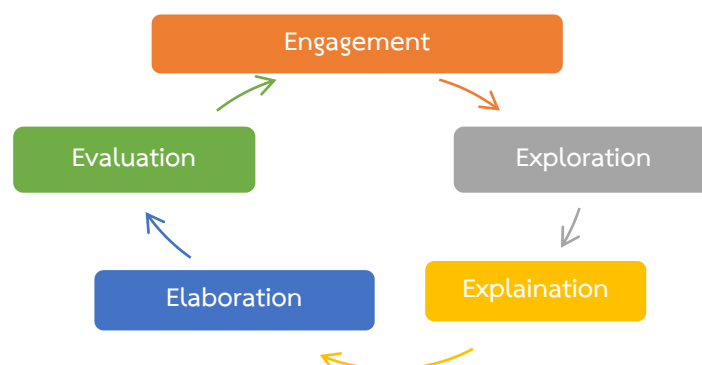


Figure 2. 5E Learning Cycle Model.

Source: Institute for the Promotion of Teaching Science and Technology (2019)

6. วิธีดำเนินงานวิจัย (Research Methodology)

6.1 ประชากร

นักเรียนชั้นมัธยมศึกษาปีที่ 1 จำนวน 138 คน โรงเรียนเทศบาล 5 สีหรัักษ์วิทยา จังหวัดอุดรธานี

6.2 กลุ่มตัวอย่าง

นักเรียนชั้นมัธยมศึกษาปีที่ 1/1 จำนวน 32 คน โรงเรียนเทศบาล 5 สีหรัักษ์วิทยา จังหวัดอุดรธานี ที่ได้จากการสุ่มแบบกลุ่ม มีขั้นตอนดังนี้ 1) ศึกษาลักษณะเบื้องต้นของประชากรแล้วจำแนกประชากรออกเป็น กลุ่มย่อยโดยที่เน้นความแตกต่างภายในกลุ่มที่แตกต่างกันคล้ายประชากร แต่จะมีความคล้ายคลึงกัน ระหว่างกลุ่มตัวอย่าง 2) สุ่มกลุ่มตัวอย่างแบบกลุ่มโดยการจับฉลากที่ระบุชื่อกลุ่มตัวอย่างแล้ว ระบุจำนวนกลุ่มตัวอย่าง

6.3 เนื้อหาที่ใช้ในการจัดการเรียนรู้

เนื้อหาในการจัดการเรียนรู้ในงานวิจัยนี้ ประกอบด้วย ความหมายของคู่อันดับ กราฟของคู่อันดับในระบบพิกัดฉาก การอ่านและแปลความหมายขงกราฟบนระนาบในระบบพิกัดฉาก และสมการและกราฟของสมการเชิงเส้นสองตัวแปร

6.4 เครื่องมือวิจัย

เครื่องมือวิจัย ประกอบด้วย แผนการจัดการเรียนรู้แบบสืบเสาะหาความรู้ แบบวัดความสามารถในการแก้ปัญหาทางคณิตศาสตร์ แบบทดสอบวัดผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ และแบบประเมินวัดความรับผิดชอบ รายละเอียดดังนี้

1) แผนการจัดการเรียนรู้แบบสืบเสาะหาความรู้ จำนวน 9 แผน แผนละ 1 ชั่วโมง รวม 9 ชั่วโมง เป็นแผนการจัดการเรียนรู้ที่เน้นให้ผู้เรียนพัฒนาความสามารถในการคิด การแก้ปัญหาหรือแสวงหาความรู้โดยใช้กระบวนการทางความคิดเพื่อหาความรู้และค้นพบคำตอบด้วยตนเอง ประกอบด้วย 5 ขั้นตอน ได้แก่ การสร้างความสนใจ การสำรวจและค้นหา การอธิบายและลงข้อสรุป การขยายความรู้ และการประเมินผล

2) แบบวัดผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ เป็นแบบปรนัย 4 ตัวเลือก จำนวน 20 ข้อ มีคุณภาพด้านความเที่ยงตรงด้วยการหาค่าความสอดคล้องระหว่างข้อคำถามกับวัตถุประสงค์ (IOC) อยู่ระหว่าง 0.67 - 1.00 ค่าความเชื่อมั่นตามวิธีการของคูเดอร์ ริชาร์ดสัน (KR-20) มีค่าเท่ากับ 0.84 ค่าความยากอยู่ระหว่าง 0.47 - 0.77 และค่าอำนาจจำแนกอยู่ระหว่าง 0.25 - 0.88

3) แบบวัดความสามารถในการแก้ปัญหาทางคณิตศาสตร์ เป็นอัตนัยจำนวน 4 ข้อ 20 คะแนน มีคุณภาพด้านความเที่ยงตรงด้วยการหาค่าความสอดคล้องระหว่างข้อคำถามกับวัตถุประสงค์ (IOC) อยู่ระหว่าง 0.67 - 1.00 ค่าความเชื่อมั่นด้วยการหาสัมประสิทธิ์แอลฟา (α -coefficient) ตามวิธีการของคอนบราค มีค่าเท่ากับ 0.80 ค่าความยากอยู่ระหว่าง 0.22 - 0.75 และค่าอำนาจจำแนกอยู่ระหว่าง 0.31 - 0.66

4) แบบสอบถามวัดความรับผิดชอบ มีลักษณะเป็นแบบสอบถามเกี่ยวกับการเอาใจใส่ต่องานที่ได้รับมอบหมาย ส่งงานตามกำหนดเวลา ตั้งใจทำงานที่ได้รับมอบหมาย จดจ่อในการทำงาน เริ่มต้นทำงานที่ได้รับมอบหมายทันที มีส่วนร่วมในการทำกิจกรรมอย่างสม่ำเสมอ มีวินัยในการทำงาน ยอมรับผลของการกระทำของตนเอง และปรับปรุงการทำงานของตนเองให้ดียิ่งขึ้น มีลักษณะเป็นแบบมาตราประเมินค่า 5 ระดับ จำนวน 9 ข้อ มีคุณภาพด้านความเที่ยงตรงด้วยการหาค่าความสอดคล้องระหว่างข้อคำถามกับวัตถุประสงค์ (IOC) อยู่ระหว่าง 0.67 - 1.00 และค่าความเชื่อมั่นด้วยการหาค่าสัมประสิทธิ์แอลฟา (α coefficient) ตามวิธีของคอนบราค มีค่าเท่ากับ 0.85

6.5 การเก็บรวบรวมข้อมูล

ดำเนินการเก็บรวบรวมข้อมูลจากนักเรียนชั้นมัธยมศึกษาปีที่ 1 โดยให้ทำแบบทดสอบวัดผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ก่อนเรียน จากนั้นดำเนินการจัดการเรียนรู้กลุ่มตัวอย่างด้วยแผนการจัดการเรียนรู้ที่สร้างขึ้นจำนวน 9 แผน โดยให้นักเรียนชั้นมัธยมศึกษาปีที่ 1 ทำกิจกรรมสร้างความสนใจด้วยการกล่าวทักทายครู กิจกรรมความหมายของคู่อันดับ และกิจกรรมวาดกราฟในระบบพิกัดฉาก กิจกรรมการคูณหาร ตามขั้นตอนการจัดการเรียนรู้แบบสืบเสาะหาความรู้ จากนั้นให้ทำแบบทดสอบวัดผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ แบบวัดความสามารถในการแก้ปัญหาทางคณิตศาสตร์ และแบบวัดความรับผิดชอบ จากนั้นนำผลที่ได้ไปวิเคราะห์ข้อมูลทางสถิติต่อไป

6.6 การวิเคราะห์ข้อมูล

ดำเนินการเปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์โดยวิเคราะห์ข้อมูลก่อนเรียนและหลังเรียน โดยใช้การทดสอบค่าที (Dependent Sample t-test) ดำเนินการเปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์กับเกณฑ์ร้อยละ 70 วิเคราะห์ข้อมูลโดยใช้การทดสอบค่าที One Sample t-test และดำเนินการศึกษาความสามารถในการแก้ปัญหาทาง

คณิตศาสตร์โดยวิเคราะห์ข้อมูลจากคะแนนรวม คำนวณค่าร้อยละและแปลผลคะแนน (Saema et al., 2021) ดังนี้ คะแนน 16 – 20 หมายถึง ระดับมาก คะแนน 11 – 15 หมายถึง ระดับปานกลาง และ คะแนน 0 – 10 หมายถึง ระดับน้อย นอกจากนี้ได้ศึกษาความรับผิดชอบหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ และวิเคราะห์ข้อมูลโดยใช้ ค่าเฉลี่ย ส่วนเบี่ยงเบนมาตรฐาน และแปลผลค่าเฉลี่ย ดังนี้ ค่าเฉลี่ย 4.50 -5.00 หมายถึง ระดับมากที่สุด ค่าเฉลี่ย 3.50 -4.49 หมายถึง ระดับมาก ค่าเฉลี่ย 2.50 -3.49 หมายถึง ระดับปานกลาง ค่าเฉลี่ย 1.50 -2.49 หมายถึง ระดับน้อย และ ค่าเฉลี่ย 1.00 -1.49 หมายถึง ระดับน้อยที่สุด

7. ผลการวิจัย (Results)

7.1 การเปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ก่อนและหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ ของนักเรียนชั้นมัธยมศึกษาปีที่ 1

ผลงานวิจัย ดัง Table 1.

Table 1. Comparing pretest and posttest Mathematics achievement of Inquiry-Based learning for grade 7 students.

Mathematics achievement	n	$\bar{X}$	S.D.	t	Sig.
Pretest	32	9.09	1.68	12.24*	.00
Posttest	32	15.75	3.54		

*p < .05

จาก Table 1. พบว่า p < .05 แสดงว่า คะแนนผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์เมื่อเรียนผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้จะมีคะแนนหลังเรียน คือ 15.75 สูงกว่าก่อนเรียน คือ 9.09 อย่างมีนัยสำคัญทางสถิติที่ระดับ .05

7.2 การเปรียบเทียบผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ หลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ ของนักเรียนชั้นมัธยมศึกษาปีที่ 1 กับเกณฑ์ร้อยละ 70

ผลงานวิจัย ดัง Table 2.

Table 2. Comparing Mathematics achievement of Inquiry-Based learning for grade 7 students with criteria of 70%.

	n	$\bar{X}$	S.D.	t	Sig.
Mathematics achievement	32	15.75	3.53	2.80*	.01

*p < .05

จาก Table 2. พบว่า p < .05 แสดงว่า คะแนนผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์หลังเรียน โดยมีค่าเฉลี่ยหลังเรียน คือ 15.75 เมื่อเรียนผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้จะมีคะแนนสูงกว่าเกณฑ์ที่กำหนดไว้ คือ ร้อยละ 70 อย่างมีนัยสำคัญทางสถิติที่ระดับ .05

7.3 การศึกษาความสามารถในการแก้ปัญหาหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ ของนักเรียนชั้นมัธยมศึกษาปีที่ 1

ผลงานวิจัย ดัง Table 3.



Table 3. Results of Mathematics problem-solving ability for grade 7 students

Mathematics Problem- Solving Ability	Scores	n	Percentage
High levels	16-20	22	68.75
Medium levels	11-15	10	31.25
Low levels	0-10	-	-
Total		32	100.00

จาก Table 3. พบว่า นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีความสามารถในการแก้ปัญหาหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ อยู่ในระดับมาก (High levels) จำนวน 22 คน คิดเป็นร้อยละ 68.75 และระดับปานกลาง (Medium levels) จำนวน 10 คน คิดเป็นร้อยละ 31.25 ตามลำดับ และสามารถแสดงผลความสามารถในการแก้ปัญหาของนักเรียนดัง Figure 3.

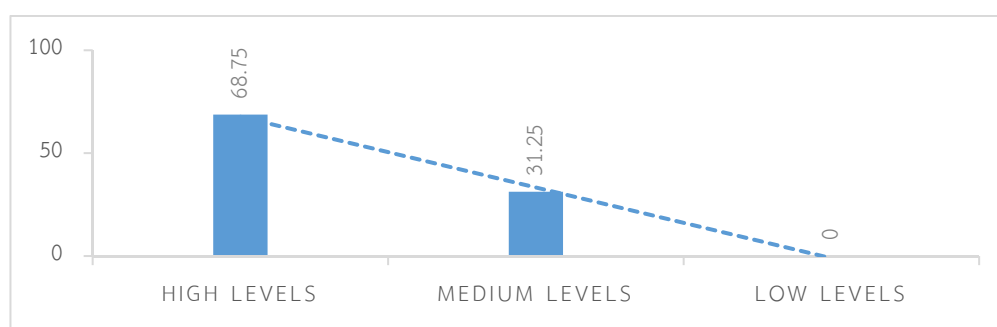


Figure 3. Mathematics problem-solving ability graph for grade 7 students

จาก Figure 3. พบว่า ร้อยละของคะแนนส่วนใหญ่อยู่ในระดับมาก ในขณะที่ร้อยละของคะแนนระดับต่ำ มีค่าเป็นศูนย์ แสดงให้เห็นว่า การจัดการเรียนรู้แบบสืบเสาะหาความรู้ของนักเรียนชั้นมัธยมศึกษาปีที่ 1 สามารถส่งเสริมความสามารถในการแก้ปัญหาของผู้เรียนได้

7.4 การศึกษาความรับผิดชอบหลังการจัดการเรียนรู้ผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้ ของนักเรียนชั้นมัธยมศึกษาปีที่ 1

ผลงานวิจัย ดัง Table 4.

Table 4. Responsibility of grade 7 students after studying inquiry-based learning.

Responsibility	$\bar{x}$	S.D.	Meaning
I pay attention to the tasks assigned to me. (ข้าพเจ้าใส่ใจต่องานที่ได้รับมอบหมาย)	4.81	0.44	Very good
I submit work on time. (ข้าพเจ้าส่งงานตามกำหนดเวลา)	4.65	0.51	Very good
I am committed to completing the tasks assigned to me. (ข้าพเจ้าตั้งใจทำงานที่ได้รับมอบหมาย)	4.42	0.68	good
I focus on my work. (ข้าพเจ้าจดจ่อในการทำงาน)	4.42	0.68	good

Responsibility	$\bar{X}$	S.D.	Meaning
I start working on assigned tasks immediately. (ข้าพเจ้าเริ่มต้นทำงานที่ได้รับมอบหมายทันที)	4.70	0.37	Very good
I participate regularly in activities. (มีส่วนร่วมในการทำกิจกรรมอย่างสม่ำเสมอ)	4.37	0.69	good
I have discipline in my work. (ข้าพเจ้ามีวินัยในการทำงาน)	4.63	0.52	Very good
I accept the consequences of my actions. (ยอมรับผลของการกระทำของตน)	4.37	0.69	good
I make improvements to my work. (ข้าพเจ้าปรับปรุงการทำงานของตนเองให้ดีขึ้น)	4.54	0.56	Very good
Total	4.55	0.57	Very good

จาก Table 4. พบว่า นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีความรับผิดชอบหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ โดยรวมอยู่ในระดับมากที่สุด ($\bar{X}$ = 4.55 และ S.D. = 0.57) เมื่อพิจารณาเป็นรายด้าน พบว่า นักเรียนเอาใจใส่ต่องานที่ได้รับมอบหมาย มีค่าเฉลี่ยสูงสุด ($\bar{X}$ = 4.81 และ S.D. = 0.44) รองลงมา คือ นักเรียนเริ่มต้นทำงานที่ได้รับมอบหมายทันที ($\bar{X}$ = 4.70 และ S.D. = 0.37) นักเรียนส่งงานตามกำหนดเวลา ($\bar{X}$ = 4.65 และ S.D. = 0.51) นักเรียนมีวินัยในการทำงาน ($\bar{X}$ = 4.63 และ S.D. = 0.52) นักเรียนปรับปรุงการทำงานของตนเองให้ดีขึ้น ($\bar{X}$ = 4.54 และ S.D. = 0.56) ตามลำดับ และค่าเฉลี่ยน้อยสุดคือ มีส่วนร่วมในการทำกิจกรรมอย่างสม่ำเสมอและยอมรับผลของการกระทำของตน ($\bar{X}$ = 4.37 และ S.D. = 0.69)

8. สรุปผลการวิจัย (Conclusion)

- ผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ของนักเรียนชั้นมัธยมศึกษาปีที่ 1 มีคะแนนเฉลี่ยหลังสูงกว่าก่อนการเรียน ด้วยการจัดการเรียนรู้แบบสืบเสาะหาความรู้ อย่างมีนัยสำคัญทางสถิติที่ระดับ.05
- ผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ของนักเรียนชั้นมัธยมศึกษาปีที่ 1 มีคะแนนเฉลี่ยหลังจากเรียนผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้ สูงกว่าเกณฑ์ที่กำหนดไว้ คือ ร้อยละ 70 อย่างมีนัยสำคัญทางสถิติที่ระดับ .05
- นักเรียนมีความสามารถในการแก้ปัญหาโดยรวมหลังจากเรียนผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้ ส่วนใหญ่อยู่ในระดับมาก จำนวน 22 คน คิดเป็นร้อยละ 68.75 และอยู่ในระดับปานกลาง จำนวน 10 คน คิดเป็นร้อยละ 31.25
- นักเรียนมีความรับผิดชอบโดยรวมหลังจากเรียนผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้ อยู่ในระดับมากที่สุด และเมื่อพิจารณาเป็นรายด้าน พบว่า นักเรียนเอาใจใส่ต่องานที่ได้รับมอบหมาย มีค่าเฉลี่ยสูงสุด รองลงมา คือ นักเรียนเริ่มต้นทำงานที่ได้รับมอบหมายทันที นักเรียนส่งงานตามกำหนดเวลา นักเรียนมีวินัยในการทำงาน และนักเรียนปรับปรุงการทำงานของตนเองให้ดีขึ้น ตามลำดับ

9. อภิปรายผลการวิจัย (Discussion)

- นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์หลังสูงกว่าก่อนการจัดการเรียนรู้แบบสืบเสาะหาความรู้ อย่างมีนัยสำคัญทางสถิติที่ระดับ .05 ซึ่งพบว่านักเรียนชั้นมัธยมศึกษาปีที่ 1 มีผลสัมฤทธิ์ทางการเรียนหลังสูงกว่าก่อนเรียน ผลดังกล่าวเนื่องจากการจัดการเรียนรู้แบบสืบเสาะหาจะช่วยพัฒนาความสามารถในการคิดของนักเรียนสามารถแก้ไขปัญหา สามารถแสวงหาความรู้ผ่านกระบวนการทางความคิดทำให้นักเรียนสามารถค้นพบคำตอบได้ด้วยตนเอง นอกจากนี้ด้วยกระบวนการเรียนการสอนที่ผู้สอนเป็นผู้สนับสนุนผู้เรียน ทำให้ผู้เรียนมีอิสระในการหาทางแก้ไข

ปัญหาด้วยตนเอง ไม่จำเป็นต้องรอคำตอบจากครู ผู้เรียนเกิดการกระตุ้นความคิดมากขึ้น จนนำไปสู่การแก้ไขปัญหาได้ ซึ่งสอดคล้องกับงานวิจัย ของ Prombungkoed et al. (2020) ได้ศึกษาความสามารถในการแก้ปัญหาทางคณิตศาสตร์โดยใช้การเรียนรู้แบบสืบเสาะหาความรู้ 5 ขั้น (5Es) ของนักเรียนชั้นมัธยมศึกษาปีที่ 3 โดยพบว่าผลสัมฤทธิ์ทางการเรียนหลังเรียนสูงกว่าก่อนเรียนอย่างมีนัยสำคัญทางสถิติที่ระดับ .05 และ Khamwari et al. (2023) พบว่า ผลสัมฤทธิ์ทางการเรียนวิชาคณิตศาสตร์โดยใช้วิธีสืบเสาะหาความรู้หลังเรียนสูงกว่าก่อนเรียนอย่างมีนัยสำคัญทางสถิติที่ระดับ .05 และเมื่อเปรียบเทียบผลสัมฤทธิ์หลังได้รับการจัดกิจกรรมการเรียนรู้ พบว่า สูงกว่าเกณฑ์ที่กำหนดไว้ คือ ร้อยละ 70 อย่างมีนัยสำคัญทางสถิติที่ระดับ .05

2. นักเรียนมีผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์หลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ สูงกว่าเกณฑ์ ร้อยละ 70 อย่างมีนัยสำคัญทางสถิติที่ระดับ.05 ผลดังกล่าวเนื่องจากการจัดการเรียนรู้แบบสืบเสาะหาความรู้ส่งเสริมให้นักเรียนพัฒนาความสามารถในการคิด การแก้ปัญหาหรือแสวงหาความรู้โดยใช้กระบวนการทางความคิดเพื่อแสวงหาความรู้และค้นพบคำตอบด้วยตนเอง นอกจากนี้ด้วยการกระตุ้นให้ผู้เรียนหาคำตอบจะทำให้ผู้เรียนพยายามเร่งหาคำตอบให้ได้ นำสู่ผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ที่สูงขึ้น (Sutthiratt, 2015) โดยมีผู้สอนเป็นผู้รู้และเสนอแนะกระตุ้นความสนใจให้ผู้เรียนเกิดความสงสัยและคิดหาคำตอบ ซึ่งสอดคล้องกับงานวิจัย Quinong et al. (2021) ได้ศึกษาผลสัมฤทธิ์ทางการเรียนวิชาคณิตศาสตร์เมื่อเรียนผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้ 5 ขั้น พบว่า สูงกว่าเกณฑ์ร้อยละ 60 อย่างมีนัยสำคัญทางสถิติที่ระดับ .05 และผลสัมฤทธิ์ทางการเรียนวิชาคณิตศาสตร์หลังเรียนสูงกว่าก่อนเรียน อย่างมีนัยสำคัญทางสถิติที่ระดับ .05

3. นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีความสามารถในการแก้ปัญหาหลังการจัดการเรียนรู้แบบสืบเสาะหาความรู้ อยู่ในระดับมาก จำนวน 22 คน คิดเป็นร้อยละ 68.75 และระดับปานกลาง จำนวน 10 คน คิดเป็นร้อยละ 31.25 ผลดังกล่าวเนื่องจากนักเรียนส่วนใหญ่ได้รับการกระตุ้นจากครูให้ค้นหาคำตอบด้วยตนเอง ตั้งแต่ขั้นตอนการวางแผนการแก้ปัญหา การลงมือแก้ไขปัญหาด้วยตัวเอง การนำเสนอผลงาน การนำไปสู่การประยุกต์ใช้งานผลที่เกิดขึ้นจากการเรียนรู้ ซึ่งเป็นไปตามขั้นตอนการเรียนรู้แบบสืบเสาะหาความรู้ ซึ่งสอดคล้องกับ (Insuwan et al., 2021) ได้พัฒนาความสามารถในการแก้ปัญหาทางคณิตศาสตร์ โดยใช้การจัดการเรียนรู้ตามแนวคิดทฤษฎี Constructivism พบว่า นักเรียนมีความสามารถในการแก้ปัญหาทางคณิตศาสตร์ หลังการจัดการเรียนรู้ โดยใช้การจัดการเรียนรู้ตามแนวคิดทฤษฎี Constructivism โดยรวมอยู่ในระดับมาก และ (Napajarin et al., 2021) ได้ศึกษาผลสัมฤทธิ์ทางการเรียนคณิตศาสตร์ การแก้ปัญหาทางคณิตศาสตร์ โดยการจัดการเรียนรู้ที่บูรณาการแนวคิด CCR พบว่า นักเรียนมีทักษะการแก้ปัญหาทางคณิตศาสตร์ หลังการจัดการเรียนรู้ที่บูรณาการแนวคิด CCR โดยรวมอยู่ในระดับ และ (Khammani, 2010) ได้ศึกษาความสามารถในการแก้ปัญหาทางคณิตศาสตร์จากการจัดการเรียนรู้โดยใช้กระบวนการเรียนรู้ 5 ขั้นตอน (5STEPS) พบว่า ความสามารถในการแก้ปัญหาทางคณิตศาสตร์ อยู่ในระดับปานกลาง

4. นักเรียนชั้นมัธยมศึกษาปีที่ 1 มีความรับผิดชอบหลังการจัดการเรียนรู้ผ่านการจัดการเรียนรู้แบบสืบเสาะหาความรู้ โดยรวมอยู่ในระดับมาก เมื่อพิจารณาข้อที่นักเรียนมีค่าเฉลี่ยมากกว่าข้ออื่น คือ ข้าพเจ้าเอาใจใส่ต่องานที่ได้รับมอบหมาย ทั้งนี้เนื่องจากผู้เรียนทราบกระบวนการเรียนรู้แบบสืบเสาะหาความรู้ ซึ่งเกิดจากครูผู้สอนได้แนะนำเบื้องต้นว่าผู้เรียนจะต้องหาคำตอบด้วยตนเอง ดังนั้นผู้เรียนจึงเกิดความเอาใจใส่ต่องานที่ได้รับมอบหมายมากขึ้น จะเห็นว่าการจัดการเรียนรู้แบบสืบเสาะหาความรู้ จะส่งเสริมให้ผู้เรียนมีความเอาใจใส่ต่อการปฏิบัติหน้าที่ที่ได้รับมอบหมาย ตั้งใจรับผิดชอบในการการทำงานให้เสร็จ มีวินัยในการทำงาน ส่งงานตามกำหนดเวลาและพัฒนาการทำงานด้วยตนเอง ยอมรับในการกระทำของตนเอง และอดทนต่ออุปสรรคในการทำงาน ซึ่งสอดคล้องกับ (Saema et al., 2021) การเรียนรู้ผ่านกิจกรรมการปฏิบัติจริง ที่ส่งเสริมความรับผิดชอบ พบว่า นักเรียนมีความรับผิดชอบหลังการจัดการเรียนรู้ผ่านกิจกรรมการปฏิบัติจริง โดยรวม



อยู่ในระดับมาก และ Sandeang et al. (2021) พบว่า นักเรียนมีความรับผิดชอบในการทำงานหลังการจัดการเรียนรู้โดยใช้การจัดการเรียนรู้ตามแนวคิดทฤษฎี Constructivism โดยรวมอยู่ในระดับมาก

10. ข้อเสนอแนะงานวิจัย (Recommendation)

10.1 ข้อเสนอแนะการนำผลการวิจัยไปใช้

1. ในการจัดการเรียนรู้ ครูผู้สอนควรตรวจสอบความรู้เดิมของนักเรียน และใช้คำถามเพื่อช่วยเหลือในการเรียนรู้ ควรคำนึงถึงความสามารถในการแก้ปัญหาของแต่ละคน เนื่องจากนักเรียนแต่ละคนความรู้พื้นฐานเดิมไม่เท่ากัน เปิดโอกาสให้นักเรียนได้แสดงความคิดเห็น และมีการเสริมแรงเพื่อกระตุ้นนักเรียนมีความสนใจในกิจกรรมการเรียนการสอน เพื่อให้การจัดกิจกรรมไปตามวัตถุประสงค์

2. สามารถนำขั้นตอนของการจัดการเรียนรู้แบบสืบเสาะหาความรู้ไปใช้ได้หลากหลายวิชา เพราะเป็นขั้นตอนที่ทำให้นักเรียนสามารถแก้ปัญหาและสร้างองค์ความรู้ได้ด้วยตนเอง แต่ควรยกตัวอย่างสถานการณ์ที่สอดคล้องกับชีวิตประจำวันหรือใกล้ตัวของนักเรียน นักเรียนสามารถทำกิจกรรมได้อย่างรวดเร็วและมั่นใจ และนักเรียนสามารถเข้าถึงสถานการณ์จริงได้

10.2 ข้อเสนอแนะการวิจัยต่อไป

1. ควรศึกษาการจัดการเรียนรู้โดยใช้การจัดการเรียนรู้แบบสืบเสาะหาความรู้ ในกลุ่มสาระการเรียนรู้คณิตศาสตร์ในระดับชั้นอื่น โดยปรับกิจกรรมการเรียนการสอน ให้เหมาะสมกับเนื้อหาในแต่ละระดับชั้น เช่น ครูให้นักเรียนทำกิจกรรมสิ่งมีชีวิตรอบตัว ครูต้องจัดกิจกรรมตามขั้นตอนของการจัดการเรียนรู้แบบสืบเสาะหาความรู้เพื่อให้นักเรียนสามารถสร้างองค์ความรู้และศึกษาหรือแก้ปัญหาได้ด้วยตนเองเพื่อให้เกิดประสิทธิภาพในการเรียนรู้ของนักเรียน

2. ควรมีการศึกษาเปรียบเทียบผลสัมฤทธิ์ทางการเรียนของวิธีการสอน โดยการจัดการเรียนการสอนแบบสืบเสาะหาความรู้ กับวิธีการสอนแบบอื่น

11. เอกสารอ้างอิง (References)

Chaiyo, J., & Sakolkeart, P. (2019). Development of Creative Thinking in Science Subject Using Inquiry Approach of Matthayomsuksa 3 Students. *Academic Journal of Humanities and Social Sciences Buriram Rajabhat University*, 11(1), 23–39. (In Thai)

Chularut, P. (2018). Learning Management for Students in the Thailand 4.0 Era. *Veridian E-Journal, Silpakorn University (Humanities, Social Sciences and Arts)*, 11(2), 2363–2380. (In Thai)

Institute for the Promotion of Teaching Science and Technology. (2019). *Supplementary Curriculum Handbook for Chemistry, Upper Secondary Level, Science Learning Area (Revised Edition 2017) under the Basic Education Core Curriculum 2008*. Chulalongkorn University Press. (In Thai)

Insuwan, S., Pantachord, U., & Chittawilai, S. (2021). The Development of Mathematics Problem Solving Ability using Learning Management on Constructivism Theory for Student Grade 9. *The 1st Kamphaeng Phet Rajabhat University Student National Conference*, 14–22.

<https://research.kpru.ac.th/research2/pages/filere/28032021-03-04.pdf>. (In Thai)

Jitchayawanit, K. (2019). *Learning Management*. Chulalongkorn University Press. (In Thai)

- Khammani, T. (2010). *The Science of Teaching: Knowledge for Effective Learning Process Organization* (13th ed.). Chulalongkorn University Press. (In Thai)
- Khamwari, S., Khansila, P., & Thienyutthakul, S. (2023). Improving Grade 8 Students' Mathematics Achievement on Factoring Quadratic Polynomials Through 5E Cycle Learning Model. *Journal of Academic Surindra Rajabhat*, 1(6), 31–44. <https://doi.org/10.14456/JASRRU.2023.38>. (In Thai)
- Napajarin Injino, N., Chaiprad, B., & Ketsri, R. (2021). Study of Mathematics Achievement Mathematic Problems Solving and Achievement Motivation by Learning Management that Integrated the CCR Concept of Polynomial for Mattayomsuksa 2. *The 1st Kamphaeng Phet Rajabhat University Student National Conference*, 62–69. <https://research.kpru.ac.th/research2/pages/filere/19702021-03-04.pdf>. (In Thai)
- Ouinong, S., Sokhuma, K., & Plangprasopchok, S. (2021). Learning Management using 5E Inquiry-Based Learning to Enhance Mathematical Communications on Pyramids, Cones and Spheres for Secondary 3 students (Grade 9). *Rajabhat Rambhai Barni Research Journal*, 15(3), 95–106. (In Thai)
- Prombungkoed, J., Boonprajak, D., & Sokhuma, K. (2020). The Effects of Learning Activities Using Inquiry Cycle (5Es) on Mathematical Problem Solving and Learning Achievement in Inequality of Mathayom Suksa 3 Students. *Rajabhat Rambhai Barni Research Journal*, 14(3), 63–71. (In Thai)
- Research Administration and Educational Quality Assurance Division. (2017). *Thailand's Development Model towards Stability, Prosperity and Sustainability*. <https://waa.inter.nstda.or.th/stks/pub/2017/20171114-draeqa-blueprint.pdf>. (In Thai)
- Saema, S., Pantachord, U., & Mansup, M. (2021). The Effect of Active Learning to Promotes Mathematical Achievement Teamwork and Responsibility of Grade 7. *The 1st Kamphaeng Phet Rajabhat University Student National Conference*, 44–52. <https://research.kpru.ac.th/research2/pages/filere/24162021-03-04.pdf>. (In Thai)
- Sandeang, K., Pantachord, U., & Sangkaew, S. (2021). The Development of Mathematics Problem Solving Ability by Using Learning Management based on Theoretical Concepts Constructivism for Student Grade 7. *The 1st Kamphaeng Phet Rajabhat University Student National Conference*, 23–32. <https://research.kpru.ac.th/research2/pages/filere/5692021-03-04.pdf>. (In Thai)
- Sinthapanon, S., Na Ranong, W., Sookying, F., Sangpirom, P., Makomon, S., Chumkot, J., Napharat, P. (2002). *Learning Process Organization: Emphasizing Learner-Centered Approach According to the Basic Education Curriculum*. Aksorn Charoen Tat. (In Thai)
- Sutthiratt, C. (2015). *80 Innovations in Learner-Centered Learning Management* (6th ed.). P Balance Design and Printing. (In Thai)



Development of Honoring His Majesty's Exhibition in Metaverse Model, Royal Rainmaking Project การพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง

Buntharic Seehabutto^{1,*}, Kamolrat Somchai², and Nitinan Mata³
บุญทฤก สีหาบุตโต^{1,*}, กมลรัตน์ สมใจ², และ นิธินันท์ มาตา³

Received: 15 March 2024;
Revised: 14 June 2024;
Accepted: 17 June 2024;
Published: 16 August 2024;

Abstract

This research aims to develop and evaluate the satisfaction of honoring His Majesty's exhibition in the metaverse model, Royal Rainmaking project in a study among 30 samples of visitors to the exhibition honoring His Majesty in data collection operations from honoring His Majesty's exhibition in the metaverse model, and satisfaction assessment form towards analysis by mean and standard deviation. Results find that viewers can avatar themselves and visit the developed metaverse exhibition space such as infographics, multimedia, and 3D models, there are comments and interactions to collect items at exhibition points. Which, is the satisfaction of honoring His Majesty's exhibition in the metaverse model of high levels.

Keywords: Honoring His Majesty's Exhibition, Metaverse Model, Royal Rainmaking Project.

บทคัดย่อ

การวิจัยนี้มีวัตถุประสงค์เพื่อพัฒนาและประเมินความพึงพอใจต่อนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวงในการศึกษากับ 30 กลุ่มตัวอย่าง คือ ผู้เข้านิทรรศการเทิดพระเกียรติในการดำเนินการเก็บข้อมูลจากนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส และแบบประเมินความพึงพอใจสู่การวิเคราะห์โดยค่าเฉลี่ย และส่วนเบี่ยงเบนมาตรฐาน ผลวิจัยพบว่า ผู้ชมสามารถ

^{1,*} Student, Program in Information Technology, Faculty of Science, Buriram Rajabhat University, Buriram 31000, Thailand; นักศึกษาสาขาวิชาเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์ มหาวิทยาลัยราชภัฏบุรีรัมย์ จังหวัดบุรีรัมย์ 31000 ประเทศไทย; Email: 630112418013@bru.ac.th

² Assistant Professor, Dr., Program in Information Technology, Faculty of Science, Buriram Rajabhat University, Buriram 31000, Thailand; ผู้ช่วยศาสตราจารย์ ดร. สาขาวิชาเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์ มหาวิทยาลัยราชภัฏบุรีรัมย์ จังหวัดบุรีรัมย์ 31000 ประเทศไทย; Email: kamonrat.sj@bru.ac.th

³ Lecturer, Dr., Program in Information Technology, Faculty of Science, Buriram Rajabhat University, Buriram 31000, Thailand; อาจารย์ ดร. สาขาวิชาเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์ มหาวิทยาลัยราชภัฏบุรีรัมย์ จังหวัดบุรีรัมย์ 31000 ประเทศไทย; Email: kamonrat.sj@bru.ac.th

*Corresponding authors: Buntharic Seehabutto (630112418013@bru.ac.th)



อวตารตนเองเข้าชมนิทรรศการพื้นที่เมตาเวิร์สที่พัฒนาขึ้น ได้แก่ สื่ออินโฟกราฟิก มัลติมีเดีย และโมเดล 3 มิติ ที่มีการแสดงความคิดเห็นและมีปฏิสัมพันธ์กับไอเทมตามจุดแสดงนิทรรศการ ซึ่งความพึงพอใจต่อนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส อยู่ในระดับมาก

คำสำคัญ: นิทรรศการเทิดพระเกียรติ, รูปแบบเมตาเวิร์ส, โครงการพระราชดำริ ฝนหลวง

1. บทนำ (Introduction)

ตลอดระยะเวลากว่า 70 ปี ที่พระบาทสมเด็จพระบรมชนกาธิเบศร มหาภูมิพลอดุลยเดชมหาราช บรมนาถบพิตร ทรงครองสิริราชสมบัติ พระองค์ได้ทรงทำงานและทรงบำเพ็ญ พระราชกรณียกิจเพื่อประโยชน์สุขของประชาชนทั้งปวง รวมทั้งได้ก่อตั้งโครงการต่าง ๆ ในพระราชดำริที่มีประโยชน์มากมายและส่งผลดีต่อประเทศชาติ ซึ่งได้เป็นแรงผลในการพัฒนาที่ยั่งยืนต่อประเทศไทยอย่างแท้จริง โดยหนึ่งในโครงการที่เหล่าชาวไทยน่าจะทราบกันดีก็คือ โครงการ “ฝนหลวง” (Innovation for Learning Division, 2019) ซึ่งเป็นหนึ่งในโครงการที่ช่วยบรรเทาปัญหาฝนแล้ง ซึ่งเกิดจากผลกระทบจากการเปลี่ยนแปลงสภาพภูมิอากาศ ทำให้เกิดพายุฤดูร้อนและพายุลูกเห็บ ดังนั้นในปี 2542 จึงเกิดโครงการปฏิบัติการฝนหลวงขึ้น โดยวิธีการทำฝนหลวงจะช่วยแก้ปัญหาฝนแล้งได้อย่างมีประสิทธิภาพในระยะสั้นของช่วงการทำเกษตรนอกฤดูของประชาชนในทุก ๆ ปีนั้นเอง (Senchai, 2018)

ดังนั้นเพื่อเป็นการสืบสานพระราชปณิธานและน้อมเกล้าน้อมกระหม่อมรำลึกในพระมหากรุณาธิคุณของพระบาทสมเด็จพระบรมชนกาธิเบศร มหาภูมิพลอดุลยเดชมหาราช บรมนาถบพิตร ทางคณะผู้วิจัยจึงนำเสนอเนื้อหาในโครงการฝนหลวงในรูปแบบนิทรรศการออนไลน์ ในรูปแบบเมตาเวิร์ส โลกเสมือนที่ถูกสร้างด้วยเทคโนโลยีดิจิทัลที่ช่วยให้ผู้ใช้สามารถเข้าร่วมและโต้ตอบกับพื้นที่เสมือน โดยนำเสนอเนื้อหาต่าง ๆ ในสภาพแวดล้อม 3 มิติที่สมจริง ทั้งข้อความ รูปภาพ ภาพกราฟิก ภาพเคลื่อนไหว เสียง คลิปวิดีโอ และโมเดล 3 มิติ ที่สามารถเปิดให้เข้าชมได้ตลอดเวลา ไม่จำกัดสถานที่ ช่วยลดค่าใช้จ่ายในการจัดสถานที่ รวมทั้งสามารถเข้าชมได้ ไม่ว่าจะอยู่ที่ใดหรือเวลาใดก็ตาม ไม่จำเป็นต้องไปถึงสถานที่จริง (Namman et al., 2019) ซึ่งจะช่วยแก้ปัญหาการจัดนิทรรศการแบบเดิม ในเรื่องของพื้นที่การนำเสนอที่จำกัด ค่าใช้จ่ายการจัดนิทรรศการที่สูงและผู้สนใจต้องเดินทางไปชมยังสถานที่นั้น ๆ นอกจากนี้ยังเป็นการสร้างประสบการณ์การชมนิทรรศการใหม่ในยุคดิจิทัล และช่วยลดปัญหาการรวมตัวกันของผู้คนซึ่งก่อให้เกิดความเสี่ยงในสถานการณ์โรคระบาดในปัจจุบันอีกด้วย

จากความสำคัญดังกล่าว คณะผู้วิจัยจึงเกิดแนวคิดที่จะพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวงขึ้น โดยมีการแบ่งพื้นที่ออกเป็น ส่วน โดยแต่ละส่วนจะมีการจัดวางองค์ประกอบต่าง ๆ เพื่อนำเสนอเนื้อหาเกี่ยวกับโครงการในพระราชดำริฝนหลวง

2. วัตถุประสงค์งานวิจัย (Research Objectives)

1. เพื่อพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง
2. เพื่อประเมินความพึงพอใจของผู้ชมที่มีต่อนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริ ฝนหลวง

3. กรอบแนวคิดงานวิจัย (Conceptual Framework)

งานวิจัยนี้กำหนดกรอบแนวคิดในการทำงานวิจัยและดำเนินการตามแนวคิดทฤษฎี ดัง Figure 1.

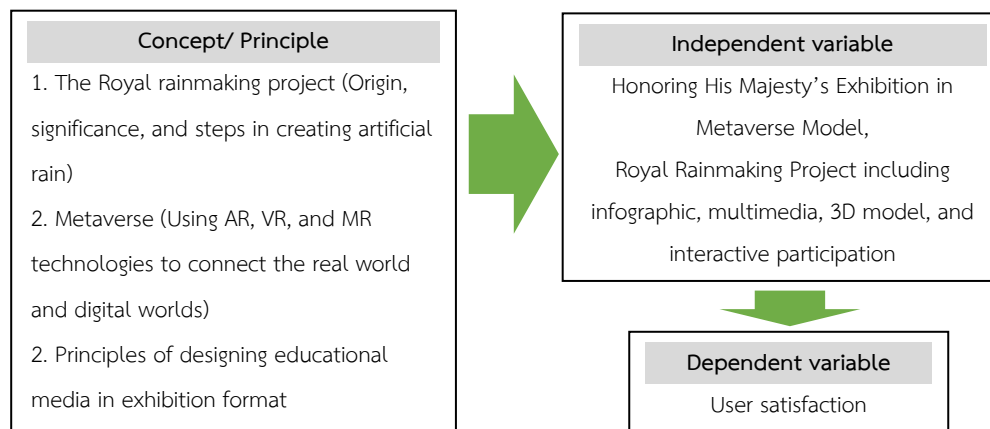


Figure 1. Conceptual Framework.

4. การทบทวนวรรณกรรมและทฤษฎีที่เกี่ยวข้อง (Literature Review)

4.1 โครงการพระราชดำริฝนหลวง

โครงการ “ฝนหลวง” เกิดขึ้นจากพระราชดำริส่วนพระองค์ของพระบาทสมเด็จพระเจ้าอยู่หัว เพื่อมีจุดประสงค์ในการสร้างฝนเทียมเพื่อบรรเทาความแห้งแล้งในภาคตะวันออกเฉียงเหนือของประเทศไทย โดยโครงการนี้เกิดจากคำรับทราบของพระบาทที่เสด็จพระราชดำเนินเยี่ยมพสกนิกรในปี พ.ศ. 2498 โดยพระองค์ทรงรับรู้ถึงสถานการณ์ความเดือดร้อนและความทุกข์ยากของราษฎร เกษตรกรที่พบปัญหาขาดแคลนน้ำ อุปโภคบริโภค และการเกษตร ดังนั้นพระมหากษัตริย์คุณพระราชทานโครงการ “ฝนหลวง” ให้ ม.ร.ว. เทพฤทธิ์ เทวกุลดำเนินการ โดยมีวัตถุประสงค์เพื่อเสริมสร้างน้ำฝนในพื้นที่ที่มีปัญหาความแห้งแล้ง (Department of Royal Rainmaking and Agricultural Aviation, 2019)

4.2 เมตาเวิร์ส (Metaverse)

โลกเสมือนที่ผู้คนสามารถเข้าไปในมิติที่สาม มีความเชื่อมโยงกับโลกความเป็นจริงและโลกดิจิทัล โดยใช้เทคโนโลยีต่าง ๆ เช่น Augmented Reality (AR) Virtual Reality (VR) และ Mixed Reality (MR) (Wuttithammapiwat et al., 2022) ซึ่งจะช่วยให้มองเห็นเป็นภาพซ้อนทับเสมือนกับในโลกของความเป็นจริง ผู้ใช้สามารถสร้างตัวละครที่สะท้อนตัวตนของตนเองในเมตาเวิร์ส และมีปฏิสัมพันธ์กับผู้คนและสิ่งของในสภาพแวดล้อมเสมือน นั้น ๆ ผ่าน Avatar หรือตัวแทนดิจิทัลของตน รวมทั้งผู้ใช้เมตาเวิร์ส สามารถเดินทางและใช้ชีวิตในสถานที่จริง สามารถมองเห็นโลกเสมือนที่เป็นโลกคู่ขนานไปพร้อมกันโดยไม่มีข้อกังวลเรื่องการเดินทางหรืออุบัติเหตุ (Sripan & Jeerapattanat, 2022)

4.3 งานวิจัยที่เกี่ยวข้อง

Marutapan et al. (2023) ได้ทำวิจัยเรื่อง การพัฒนารูปแบบการสื่อสารผ่านเมตาเวิร์สเพื่อส่งเสริมการเข้าถึงการให้บริการสารสนเทศของสำนักคอมพิวเตอร์และเครือข่าย โดยมีวัตถุประสงค์ เพื่อพัฒนาบริการสารสนเทศ และประชาสัมพันธ์บริการของสำนักคอมพิวเตอร์และเครือข่ายผ่าน เมตาเวิร์ส กลุ่มตัวอย่างเป็นผู้ตอบแบบสอบถามจำนวน 31 คน พบว่า ผลประเมินความพึงพอใจในด้านการรับรู้การให้บริการ มีค่ามากขึ้น หลังการเข้าชมพื้นที่เมตาเวิร์สของสำนัก

คอมพิวเตอร์และเครือข่าย ผู้ตอบแบบสอบถามรู้จักสำนักคอมพิวเตอร์ และเครือข่ายมากขึ้น โดยมีระดับความพึงพอใจเฉลี่ย ก่อนเข้าชมและหลังเข้าชมเพิ่มจาก 3.90 เป็น 4.48 ซึ่งมีค่าอยู่ใน ระดับมาก

Boonchuen et al. (2023) ได้ทำวิจัยเรื่องการพัฒนาสื่อการเรียนรู้ 3 มิติผ่านเมตาเวิร์ส กรณีศึกษาคลองแม่ข่า โดยมีวัตถุประสงค์ เพื่อพัฒนาเมตาเวิร์ส คลองแม่ข่าและเพื่อศึกษาความพึงพอใจจากการใช้งานเมตาเวิร์ส คลองแม่ข่า ผลการวิจัยพบว่า เมตาเวิร์ส คลองแม่ข่า สามารถนำมาใช้เพื่อประกอบการสอนในรายวิชาวิทยาการคำนวณ ร่วมกับการเรียนการสอนแบบระยะไกลหรือออนไลน์ได้เป็นอย่างดี และผู้ใช้งานมีความพึงพอใจอยู่ในระดับมาก ค่าเฉลี่ยเท่ากับ 4.46

Supat (2023) ได้ทำวิจัยเรื่อง “The Author Room” นิทรรศการเสมือนจริง บนแพลตฟอร์ม Spatial โดยมีวัตถุประสงค์เพื่อพัฒนาและประเมินประสิทธิภาพนิทรรศการเสมือนจริง “The Author Room” บนแพลตฟอร์ม Spatial และเพื่อประเมินความพึงพอใจต่อนิทรรศการเสมือนจริง “The Author Room” บนแพลตฟอร์ม Spatial ผลการวิจัยพบว่าประสิทธิภาพด้านเนื้อหา ด้านการออกแบบและด้านการใช้งาน อยู่ในระดับมากที่สุด 4.71 สามารถนำไปใช้งานจริงได้ โดยการออกบูธจัดแสดงนิทรรศการเสมือนจริงผ่านอุปกรณ์แว่น VR ในงานประชุมวิชาการทั้งระดับชาติและนานาชาติ ผลสำรวจความพึงพอใจภาพรวมต่อการรับชมค่าเฉลี่ย 4.50 อยู่ในเกณฑ์ระดับมากที่สุด

5. วิธีดำเนินงานวิจัย (Research Methodology)

5.1 ขั้นตอนการวิเคราะห์ (Analysis)

1. วิเคราะห์เนื้อหา โครงการ “ฝนหลวง” โดยทำการรวบรวมเนื้อหา แยกแยะเนื้อหา จัดลำดับเนื้อหา รวมทั้งดูความเหมาะสมกับทรัพยากรและข้อจำกัดในการแสดงเนื้อหา

2. วิเคราะห์กลุ่มเป้าหมาย ได้แก่บุคคลทั่วไปที่มีความสนใจต้องการเข้าเยี่ยมชมนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส

3. วิเคราะห์แนวคิดการพัฒนานิทรรศการพื้นที่เมตาเวิร์ส โดยนำเสนอเนื้อหาโครงการฝนหลวงด้วยสื่ออินโฟกราฟิก มัลติมีเดียและโมเดล 3 มิติ โดยผู้ชมสามารถเปิดให้เข้าชมได้ตลอดเวลา ไม่จำกัดสถานที่ และยังช่วยลดปัญหาการรวมตัวกันของผู้คน

5.2 ขั้นตอนการออกแบบ (Design)

การออกแบบมีขั้นตอนการออกแบบดัง Figure 2.

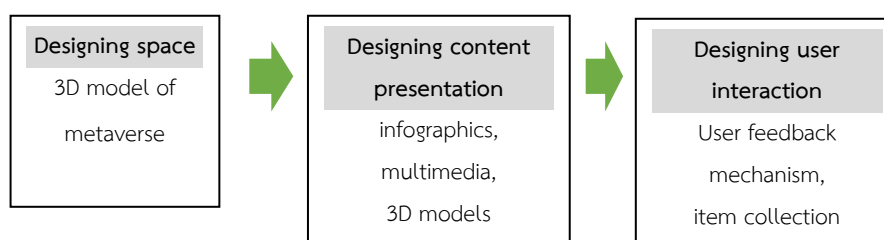


Figure 2. Design process.

1. การออกแบบพื้นที่ (Space) โมเดล 3 มิติเป็นพื้นที่เมตาเวิร์สแสดงนิทรรศการ จุดการเข้าชมตามแผนผังการจัดนิทรรศการ ดัง Figure 3.

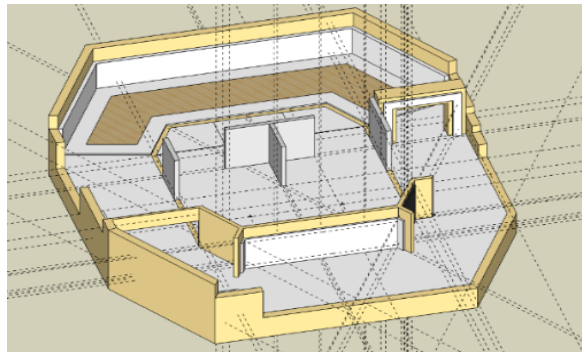


Figure 3. Isometric view of the exhibition space.

2. การออกแบบการนำเสนอเนื้อหาประกอบด้วย 1) สื่ออินโฟกราฟิก จำนวน 8 แผ่น 2) มัลติมีเดียแสดงขั้นตอนการสร้างฝนหลวง 3) โมเดลจำลองเครื่องบินการทำฝนหลวง 4) โมเดลประกอบฉากต่าง ๆ และ 5) รายการสิ่งของ (ไอเทม)

3. การออกแบบปฏิสัมพันธ์ของผู้เข้าชมในลักษณะผู้ชมกับเทคโนโลยี โดยการแสดงความคิดเห็น และเก็บรายการสิ่งของ (ไอเทม) ที่เตรียมไว้ตามจุดแสดงนิทรรศการ

5.3 ขั้นตอนการพัฒนา (Development)

สำหรับขั้นตอนในการพัฒนาคณะผู้วิจัยได้พัฒนาโดยใช้แพลตฟอร์ม Spatial ซึ่งมีขั้นตอนดังนี้

1. สร้างพื้นที่เมตาเวิร์สในสภาพแวดล้อม 3 มิติ ห้องแสดงนิทรรศการเทิดพระเกียรติ โครงการพระราชดำริฝนหลวง โดยใช้โปรแกรม Sketch up

2. สร้างเนื้อหาแนะนำประกอบด้วย 1) สร้างโมเดล 3 มิติเป็นแบบจำลองเครื่องบินในการทำฝนหลวง โมเดลประกอบฉากต่าง ๆ ด้วยโปรแกรม Blender 2) สร้างสื่ออินโฟกราฟิก 3) สร้างมัลติมีเดียแสดงขั้นตอนการสร้างฝนหลวง และ สร้างรายการสิ่งของ (ไอเทม)

3. จัดวางองค์ประกอบต่าง ๆ ได้แก่ พื้นที่เมตาเวิร์ส เนื้อหา และรายการสิ่งของต่าง ๆ ลงในโปรแกรม Unity และนำเข้าทดสอบในแพลตฟอร์ม Spatial

5.4 ขั้นตอนการนำไปใช้ (Implement)

สำหรับนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง คณะผู้วิจัยได้นำไปให้ผู้สนใจได้เข้าไปเยี่ยมชมนิทรรศการได้ที่ <https://shorturl.asia/ZY1ww> และทำการเก็บรวบรวมข้อมูลจากจุดเชื่อมต่อลิงค์เข้าสู่ Google Forms เพื่อประเมินความพึงพอใจทาง จากกลุ่มตัวอย่าง จำนวน 30 คน โดยการสุ่มอย่างง่าย แบบเจาะจง (Purposive Sampling)

5.5 ขั้นตอนการประเมินผล (Evaluation)

สำหรับการประเมินผลคณะผู้วิจัยได้ทำแบบประเมินคุณภาพโดยผู้เชี่ยวชาญ และแบบสอบถามประเมินความพึงพอใจของผู้เข้าชม แบบมาตราส่วนประมาณค่า (Rating Scale) 5 ระดับ โดยใช้สถิติที่ในการวิเคราะห์ข้อมูลได้แก่ ค่าเฉลี่ย และ ส่วนเบี่ยงเบนมาตรฐาน ไปเปรียบเทียบกับเกณฑ์การประเมิน (Arreerard, 2010) ดังนี้

ค่าเฉลี่ย 4.50-5.00 หมายถึง พึงพอใจมากที่สุด

ค่าเฉลี่ย 3.50 - 4.49 หมายถึง พึงพอใจมาก

ค่าเฉลี่ย 2.50 - 3.49 หมายถึง พึงพอใจปานกลาง

ค่าเฉลี่ย 1.50 - 2.49 หมายถึง พึงพอใจน้อย

ค่าเฉลี่ย 1.00 - 1.49 หมายถึง พึงพอใจน้อยที่สุด

6. ผลการวิจัย (Results)

ผลการดำเนินการวิจัยการพัฒนาพิพิธภัณฑ์เทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง สามารถเข้าถึงได้ที่ <https://shorturl.asia/ZY1ww> และมีผลการพัฒนาตามวัตถุประสงค์การวิจัยดังต่อไปนี้

6.1 ผลการพัฒนาพิพิธภัณฑ์เทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง

พิพิธภัณฑ์เทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการ พระราชดำริฝนหลวง มีพัฒนาโดยจัดแสดงนิทรรศการ ออกเป็นโซนการแสดง ได้แก่ โซนที่ 1 จุดเริ่มต้นแสดงแผนผังการเข้าชมนิทรรศการ โซนที่ 2 จุดแสดงด้านหน้าพิพิธภัณฑ์ โซนที่ 3 จุดแสดงสื่ออินโฟกราฟิกเนื้อหาของโครงการฝนหลวงและมัลติมีเดีย โซนที่ 4 จุดแสดงโมเดลแบบจำลอง และโซนที่ 5 จุดเก็บรายการสิ่งของ ดัง Figure 4. ถึง Figure 9.



Figure 4. Aerial view of the space in the metaverse.



Figure 5. Zone 1: starting point and displaying the visitor map



Figure 6. Zone 2: displaying an area in front of the exhibition.

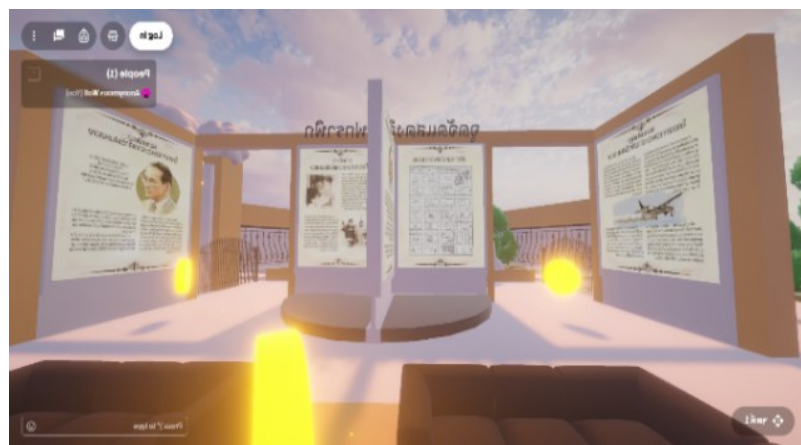


Figure 7. Zone 3: showing infographics to present content of the Royal Rainmaking Project, and multimedia.



Figure 8. Zone 4: showing simulation of models.

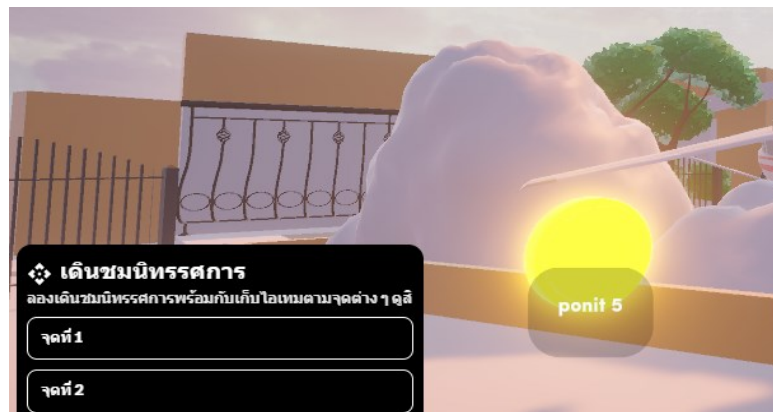


Figure 9. Zone 5: displaying the storage point for items.

6.2 ผลการประเมินความพึงพอใจของผู้เข้าเยี่ยมชมนิทรรศการ

ผลการประเมินความพึงพอใจของผู้เข้าเยี่ยมชมนิทรรศการ จากกลุ่มตัวอย่างจำนวน 30 คน สามารถแสดงผลได้ดัง Figure 10.

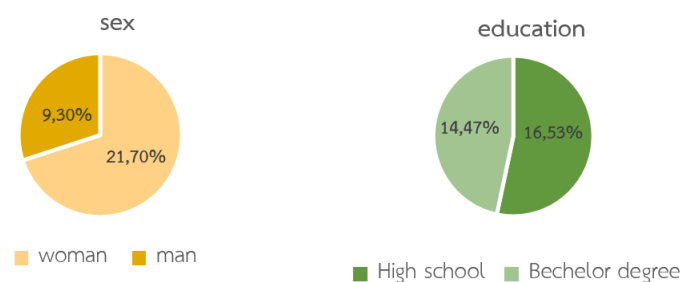


Figure 10. Description of sampling classified by sex and education.

จาก Figure 10. แสดงร้อยละของผู้เข้าชมนิทรรศการพบว่า ผู้เข้าชมนิทรรศการเป็นหญิง ร้อยละ 70 โดยผู้เข้าชมส่วนมากอยู่ในระดับการศึกษาปริญญาตรี ร้อยละ 53

ผลการประเมินความพึงพอใจของผู้เข้าชมนิทรรศการ พบว่า ผลการประเมินความพึงพอใจ ของผู้เข้าชมนิทรรศการ พบว่ามีความพึงพอใจอยู่ในระดับมาก มีค่าเฉลี่ยเท่ากับ 4.41 ± 0.37 มีพิจารณาารายข้อพบว่า ผู้ชมมีความพึงพอใจ การนำเสนอสื่อที่ชัดเจน สามารถสื่อสารให้ผู้ชมเข้าใจง่าย มีค่าเฉลี่ยเท่ากับ 4.46 ± 0.67

7. สรุปผลการวิจัย (Conclusion)

งานวิจัยมีวัตถุประสงค์เพื่อพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง และเพื่อประเมินความพึงพอใจของผู้ชมที่มีต่อนิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง สามารถสรุปผลการวิจัย ดังนี้ นิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส โครงการพระราชดำริฝนหลวง แบ่งออกเป็น 5 โซน ได้แก่ โซนที่ 1 จุดเริ่มต้นแสดงแผนผังการเข้าชมนิทรรศการ โซนที่ 2 จุดแสดงด้านหน้านิทรรศการ โซนที่ 3 จุดแสดงสื่ออินโฟกราฟิกเนื้อหาของโครงการฝนหลวงและมัลติมีเดีย โซนที่ 4 จุดแสดงโมเดลแบบจำลอง และโซนที่ 5 จุดเก็บ

รายการสิ่งของ มืองค์ประกอบของสื่อเมตาเวิร์ส ได้แก่ อินโฟกราฟิก สื่อประสม โมเดลสามมิติ และการปฏิสัมพันธ์กับ
ผู้ใช้งาน และผลการประเมินความพึงพอใจ ของผู้เข้าชมนิทรรศการพบว่ามีความพึงพอใจอยู่ในระดับมาก

8. อภิปรายผลการวิจัย (Discussion)

การพัฒนานิทรรศการเทิดพระเกียรติในรูปแบบเมตาเวิร์ส คณะผู้วิจัยได้ดำเนินการ โดยการประยุกต์ใช้ ADDIE Model เป็นหลักการออกแบบสื่อและพัฒนาระบบ เนื่องจากการศึกษาพบว่า ADDIE Model สามารถส่งผลให้สื่อมีประสิทธิภาพ เกิดการเรียนรู้หลากหลาย (Duangnim & Kareng, 2023) โดยมีกระบวนการพัฒนาได้ 5 ขั้นตอน ได้แก่ การวิเคราะห์ (Analysis) การออกแบบ (Design) การพัฒนา (Development) ได้แก่การสร้างองค์ประกอบต่าง ๆ เช่น ออกแบบพื้นที่ (space) สร้างโมเดล 3 มิติ สร้างเนื้อหาในรูปสื่ออินโฟกราฟิกและมัลติมีเดีย สร้างรายการสิ่งของ (ไอเท็ม) ตามแนวคิดของการสร้างสภาพแวดล้อมที่ผู้ใช้สามารถโต้ตอบได้ในรูปแบบ 3 มิติ แล้วรวบรวมองค์ประกอบทั้งหมดด้วย โปรแกรม Unity และนำเข้าทดสอบในแพลตฟอร์ม Spatial การนำไปใช้ (Implement) และการประเมินผล (Evaluation) ซึ่งผลการประเมินความพึงพอใจ จากกลุ่มตัวอย่างผู้เข้าชม 30 คน พบว่ามีความพึงพอใจในระดับมาก โดยเฉพาะด้านการนำเสนอเนื้อหาที่ชัดเจนและเข้าใจง่าย ซึ่งสะท้อนถึงความสำเร็จในการใช้สื่ออินโฟกราฟิกและมัลติมีเดียในการสื่อสารเนื้อหา ผู้ชมส่วนใหญ่เป็นผู้หญิงและมีระดับการศึกษาปริญญาตรี แสดงให้เห็นว่านิทรรศการนี้สามารถดึงดูดความสนใจจากกลุ่มเป้าหมายที่หลากหลายได้ สอดคล้องกับ Supat (2023) ที่ได้ทำวิจัยเรื่อง “The Author Room” นิทรรศการเสมือนจริง บนแพลตฟอร์ม Spatial และพบว่าผลสำรวจความพึงพอใจภาพรวมต่อการรับชม อยู่ในระดับมากที่สุด สอดคล้องกับงานวิจัยของ Marutapan et al. (2023) ที่ได้ทำการศึกษาวิจัยเรื่องการพัฒนาแบบการสื่อสารผ่านเมตาเวิร์สเพื่อส่งเสริมการเข้าถึงการให้บริการสารสนเทศของสำนักคอมพิวเตอร์และเครือข่าย พบว่าผู้ใช้บริการมีความพึงพอใจเพิ่มขึ้นจากการใช้งานเมตาเวิร์ส โดยระดับความพึงพอใจเฉลี่ยเพิ่มจาก 3.90 เป็น 4.48 หลังการเข้าชม ซึ่งแสดงให้เห็นถึงประสิทธิภาพของเมตาเวิร์สในการส่งเสริมการสื่อสารและการรับรู้ข้อมูล และนอกจากนี้ Boonchuen et al. (2023) ได้ทำวิจัยเรื่องการพัฒนาสื่อการเรียนรู้ 3 มิติผ่านเมตาเวิร์ส กรณีศึกษาคลองแม่ข่า ซึ่งพบว่าเมตาเวิร์สสามารถนำมาใช้เป็นสื่อการสอนร่วมกับการเรียนการสอนออนไลน์ได้อย่างดี ผู้ใช้งานมีความพึงพอใจอยู่ในระดับมาก ค่าเฉลี่ยเท่ากับ 4.46 ซึ่งงานวิจัยดังกล่าวเน้นย้ำถึงความสำเร็จในการใช้เทคโนโลยีเมตาเวิร์สเพื่อเพิ่มประสิทธิภาพในการสื่อสารและการเรียนรู้ นอกจากนี้ในด้านการประยุกต์กับการจัดเก็บข้อมูลท้องถิ่นชุมชนในรูปแบบของสารสนเทศที่มีความจำเป็นต้องนำเสนอให้ผู้ชมเข้าถึงได้โดยง่ายนั้น สอดคล้องกับงานวิจัยของ Kewsavang & Supadit (2023) ที่มีการสร้างสภาพแวดล้อมในรูปแบบสามมิติ โดยผ่านการจำลองวัตถุจริงและจัดทำให้อยู่ในรูปแบบวัตถุสามมิติ และเข้าถึงโดยง่ายผ่านเทคโนโลยีที่ทันสมัยทั้งในรูปแบบ AR, VR และเมตาเวิร์ส

9. ข้อเสนอแนะงานวิจัย (Recommendation)

ข้อเสนอแนะสำหรับการทำวิจัยครั้งต่อไป มีดังนี้ คือ 1) สร้างความรู้สึกให้ผู้เข้าชมรู้สึกเหมือนมีส่วนร่วมกับนิทรรศการให้มากขึ้น โดยการเพิ่มเกม การกิจ หรือกิจกรรมอื่น ๆ ที่หลากหลาย และ 2) เพิ่มการเชื่อมต่อไปยัง Space อื่น ๆ ที่มีเนื้อหาเกี่ยวกับโครงการเนื่องในพระราชดำริ เป็นต้น

10. เอกสารอ้างอิง (References)

Arreerard, P. (2010). *Educational Software Development*. Faculty of Education, Rajabhat Maha Sarakham University. (In Thai)



- Boonchuen, C., Chueanongkwai, T., & Kankhat, S. (2023). Development of 3D Learning Materials through Metaverse Mae Kha Canal Case Study. *Science and Technology to Community*, 1(2), 27–35.
<https://doi.org/10.57260/stc.2023.536>. (In Thai)
- Department of Royal Rainmaking and Agricultural Aviation. (2019). *History of the Department*.
https://www.royalrain.go.th/royalrain/Editor_Page.aspx?MenuId=5. (In Thai)
- Duangnim, A., & Kareng, A. (2023). Development of the Metaverse Classroom of the Division of Technology and Learning Innovation, Office of Academic Resources, Prince of Songkla University Pattani Campus. *The 13th PULINET National Conference PULINET 2023 : Foresight in Reinventing Libraries in a VUCA World : The Next Move*, 111-121.
- Innovation for Learning Division. (2019). *Video Production Guide - Innovation for Learning Division*.
<https://ita.ipst.ac.th/wp-content/uploads/sites/77/2019/06/คู่มือการผลิตวีดิทัศน์ผ่านนวัตกรรมเพื่อการเรียนรู้.pdf>. (In Thai)
- Kewsavang, A., & Supadit, K. (2023). Experience in Deveopling Local Wisdom of Nontaburi Province by Using AR, VR Technology and Metaverse. *The 13th PULINET National Conference PULINET 2023 : Foresight in Reinventing Libraries in a VUCA World : The Next Move*, 87-98. (In Thai)
- Marutapan, N., Wanram, S., & Suriya, A. (2023). Development of Communication Model on Metaverse for Access to Information Services of the Office of Computer and Networking. *17th Ubon Ratchathani University Research Conference: Research and Innovation in a Changing World*, 243–252. (In Thai)
- Namman, C., Boonsuya, A., & Thongsiri, W. (2019). The Development of Virtual Exhibitions System using Google Cardboard. *Engineering Journal of Siam University*, 20(1), 2–11. (In Thai)
- Senchai, N. (2018). Royal Rain Volunteers Network with Increasing Efficiency In Royal Rain Services. *Mahachula Academic Journal*, 5(2), 79–90. (In Thai)
- Sripan, T., & Jeerapattanatorn, P. (2022). Exnovation: The Ultimate Development of Innovation and Roles of Metaverse in Education and Training in the Next Normal Era. *Journal of Innovation and Management*, 7(2), 174–188. (In Thai)
- Supat, N. (2023). *"The Author Room" Virtual Exhibition on the Spatial Platform*.
https://www.royalrain.go.th/royalrain/Editor_Page.aspx?MenuId=5. (In Thai)
- Wuttithammapiwat, S., Salee, P., & Sainago, C. (2022). Academic Article On the Teaching Management of Thai Literature Studies Using Metaverse. *Journal of Dhammasuksa Research*, 5(2), 282–303. (In Thai)