

การใช้เทคนิคเหมืองข้อมูลในการพัฒนาโมเดลสำหรับทำนายสาขาวิชา ที่เหมาะสมกับผู้สมัครเรียน มหาวิทยาลัยราชภัฏกำแพงเพชร

Applying Data Mining Techniques to Develop a Model for Predicting Suitable Academic Fields for University Applicants at Kamphaeng Phet Rajabhat University

กนกวรรณ เขียววัน^{1*}, นรุตม์ บุตรพลอย¹, จินดาพร อ่อนเกต², ฆัมภิกา ดันติสันติสม²,

พรหมเมศ วีระพันธ์³ และ คมกริช กลิ่นอาจ³

Kanokwan Khiewwan^{1*}, Narut Butploy¹, Jindaporn Ongate², Khumphicha Tantisantisom²,

Phrommate Veeraphan³ and Komkrit Klin-art

¹สาขาวิชาเทคโนโลยีคอมพิวเตอร์ มหาวิทยาลัยราชภัฏกำแพงเพชร

²สาขาวิชาเทคโนโลยีสารสนเทศ มหาวิทยาลัยราชภัฏกำแพงเพชร

³สำนักส่งเสริมวิชาการและงานทะเบียน มหาวิทยาลัยราชภัฏกำแพงเพชร

69 หมู่ 1 ต.นครชุม อ.เมือง จ.กำแพงเพชร 62000

ผู้นิพนธ์ประสานงาน : kanokwan_kh@kpru.ac.th

วันที่รับบทความ: 26 ธันวาคม 2567 / วันที่แก้ไขบทความ ครั้งที่ 1: 12 มีนาคม 2568 / วันที่ตอบรับการตีพิมพ์: 8 พฤษภาคม 2568

วันที่แก้ไขบทความ ครั้งที่ 2: 26 มีนาคม 2568

บทคัดย่อ การวิจัยนี้มีวัตถุประสงค์เพื่อพัฒนาโมเดลสำหรับการทำนายความเหมาะสมของสาขาวิชาที่นักศึกษาในระดับปริญญาตรีเลือกเรียน โดยประยุกต์ใช้เทคนิคเหมืองข้อมูลจำนวน 5 โมเดล ได้แก่ ต้นไม้ตัดสินใจ การเรียนรู้แบบอย่างง่าย ซัพพอร์ตเวกเตอร์แมชชีน แรนดอมฟอเรสต์ และอะดามัสต์ บนชุดข้อมูลของนักศึกษาจำนวน 1,392 คนจาก 10 หลักสูตรของมหาวิทยาลัยราชภัฏกำแพงเพชร โดยจัดกลุ่มระดับผลสัมฤทธิ์ทางการเรียนเป็น 4 กลุ่ม ได้แก่ ดีมาก (A), ดี (B), ปานกลาง (C) และ ไม่ผ่าน (E) ผลการทดลองพบว่า การเรียนรู้แบบอย่างง่ายให้ค่าเฉลี่ย F1-Score สูงสุดใน 6 จาก 10 หลักสูตร โดยเฉพาะในกลุ่มสาขาวิชาที่มีลักษณะข้อมูลเรียบง่ายและแยกกลุ่มได้ชัดเจน ขณะที่ซัพพอร์ตเวกเตอร์แมชชีนแสดงประสิทธิภาพดีในหลักสูตรที่มีโครงสร้างข้อมูลซับซ้อน และแรนดอมฟอเรสต์มีความสามารถโดดเด่นในการจำแนกข้อมูลที่มีความแปรปรวนสูง โดยเฉพาะในหลักสูตรการจัดการทั่วไป ซึ่งให้ค่า F1-Score สูงสุดที่ 0.75 ผลการวิเคราะห์แสดงให้เห็นว่า การเลือกโมเดลที่เหมาะสมควรพิจารณาพร้อมกับโครงสร้างของข้อมูลในแต่ละบริบทด้วย ทั้งนี้แม้ว่า การเรียนรู้แบบอย่างง่ายจะให้ผลลัพธ์โดยรวมดีที่สุด แต่ลักษณะของข้อมูลที่มีการซ้อนทับกันระหว่างคลาสในหลายหลักสูตร ยังคงเป็นข้อจำกัดที่ทำให้ระดับความแม่นยำของการจำแนกอยู่ในระดับปานกลาง ข้อเสนอแนะสำหรับการวิจัยในอนาคต คือ การพัฒนาโมเดลที่รวมข้อมูลเชิงพฤติกรรม ความสนใจ หรือทักษะด้านอื่นของผู้เรียน เพื่อเพิ่มประสิทธิภาพในการทำนาย และสนับสนุนระบบแนะแนวการศึกษาที่สอดคล้องกับศักยภาพของนักศึกษาได้อย่างเหมาะสม

คำสำคัญ: เทคนิคเหมืองข้อมูล, การทำนายสาขาวิชา, มหาวิทยาลัยราชภัฏกำแพงเพชร, ระดับผลสัมฤทธิ์ทางการเรียน

Abstract This research aimed to develop predictive models for determining the suitability of academic programs for undergraduate students by applying five data mining techniques: Decision Tree (DT), Naive Bayes, Support Vector Machine (SVM), Random Forest, and AdaBoost. The study utilized a dataset of 1,392 students across 10 academic programs at Kamphaeng Phet Rajabhat University. Students' academic performance was grouped into four levels include Excellent (A), Good (B), Fair (C), and Low (E). The experimental results indicated that Naive Bayes achieved the highest average F1-Score in 6 out of 10 programs, particularly in programs with simple data structures and clearly separable classes. In contrast, SVM performed well in programs with complex or overlapping data structures, while Random Forest demonstrated outstanding performance in handling high-variance data, especially in the General Management program, where it achieved the highest F1-Score of 0.75. The findings suggest that selecting an appropriate model should consider the underlying structure of the dataset in each specific context. Although Naive Bayes yielded the best overall results, data overlapping between classes in several programs remained a limiting factor, resulting in moderate classification accuracy. Future research should consider incorporating behavioral, interest-based, or skill-related features to enhance prediction accuracy and support educational guidance systems that better align with each student's potential.

Keywords: Data Mining, Program Prediction, Kamphaeng Phet Rajabhat University, Students' academic performance

1. บทนำ

การเลือกสาขาวิชาที่เหมาะสมสำหรับนักศึกษาปี 1 เป็นหนึ่งในปัจจัยสำคัญที่มีผลกระทบต่อความสำเร็จทางการศึกษาและการพัฒนาอาชีพในอนาคต อย่างไรก็ตาม นักเรียนชั้นมัธยมศึกษาชั้นปีที่ 6 จำนวนมาก ประสบปัญหาการตัดสินใจเลือกสาขาวิชาเนื่องจากขาดข้อมูลที่ชัดเจนเกี่ยวกับความถนัด ความเหมาะสม และโอกาสในตลาดแรงงาน การตัดสินใจที่ผิดพลาดอาจทำให้เกิดปัญหานักศึกษาออกกลางคัน เพราะไม่มีความสนใจและกระตือรือร้นในการเรียน กรณีของมหาวิทยาลัยราชภัฏกำแพงเพชรพบปัญหาอัตราการออกกลางคันสูงถึงร้อยละ 25 ในแต่ละปีการศึกษา [1] โดยปัญหาสำคัญเกิดจากนักเรียนขาดประสบการณ์ ไม่ทราบถึงความต้องการและทักษะของตน ประกอบกับไม่รู้จักแต่ละสาขาวิชามากพอ ทำให้ใช้ความรู้สึกชอบ มองแค่สภาพแวดล้อม ตามเพื่อน หรือความเห็นของผู้ปกครองในการเลือกสาขาวิชา [2] ซึ่ง

สะท้อนถึงความจำเป็นในการพัฒนาระบบการแนะนำที่มีประสิทธิภาพ

เทคโนโลยีการขุดค้นข้อมูล (Data Mining) เป็นเครื่องมือสำคัญที่ช่วยในการวิเคราะห์และประมวลผลข้อมูลจำนวนมากเพื่อค้นหารูปแบบและความสัมพันธ์ที่ซ่อนอยู่ในข้อมูล [3][4] งานวิจัยการใช้เทคนิคการขุดค้นข้อมูลในงานด้านการศึกษา [5] นำเสนอการใช้เทคนิคการจำแนกประเภท (Classification) เพื่อทำนายผลการเรียนของนักศึกษา โดยอาศัยข้อมูลประชากรและผลการเรียนในอดีต พบว่าอัลกอริทึม Decision Tree และ Naive Bayes มีความแม่นยำสูงในการทำนายผลสัมฤทธิ์ทางการศึกษา เทคนิคเหล่านี้สามารถประยุกต์ใช้ในระบบแนะนำการศึกษาด้วยการวิเคราะห์ข้อมูลเชิงลึกในระดับสากล งานวิจัยของ [6] ชี้ให้เห็นว่าเทคนิคเหมืองข้อมูล เช่น Decision Tree และ Naive Bayes ได้ถูกนำไปใช้ในการศึกษาระบบแนะนำและการพยากรณ์ผลสัมฤทธิ์ทางการศึกษาของนักเรียนในหลายประเทศ โดยพบว่า

อัลกอริทึมเหล่านี้สามารถช่วยในการจำแนกนักเรียนที่มีความเสี่ยงต่อการออกกลางคันได้อย่างแม่นยำ งานวิจัยของ [7] เปรียบเทียบวิธีการทำเหมืองข้อมูลต่าง ๆ โดยใช้ข้อมูลจากการสำรวจนักศึกษาชั้นปีที่ 1 ของคณะเศรษฐศาสตร์ มหาวิทยาลัย Tuzla ในปีการศึกษา 2010-2011 ผลการศึกษาพบว่าตัวแปรทางสังคมและประชากรของนักศึกษา ผลการเรียนรู้จากโรงเรียนมัธยม และทัศนคติต่อการเรียนมีผลต่อความสำเร็จทางการศึกษา ซึ่งสอดคล้องกับงานวิจัยของ [8] ที่พบว่าปัจจัยทางพฤติกรรม เช่น ความสม่ำเสมอในการเข้าเรียนและการมีส่วนร่วมในชั้นเรียน มีผลต่อโอกาสประสบความสำเร็จในการศึกษา งานวิจัยการจำแนกประเภทด้วย Decision Tree ดังเช่น [9] พัฒนาอัลกอริทึม Decision Tree ที่ชื่อว่า ID3 และ C4.5 ซึ่งใช้สำหรับจำแนกข้อมูลและสร้างโมเดลทำนายที่มีความแม่นยำ งานวิจัยนี้เป็นจุดเริ่มต้นของการพัฒนาอัลกอริทึม Decision Tree ที่เข้าใจง่ายและใช้งานได้หลากหลาย เช่น การคัดเลือกสาขาวิชาและพยากรณ์ความสำเร็จในด้านการศึกษา งานวิจัยของ [10] ศึกษาการเลือกสาขาวิชาของนักศึกษาระดับปริญญาตรีโดยใช้เทคนิค Decision Tree และ Naive Bayes พบว่า Decision Tree มีความสามารถในการวิเคราะห์และทำนายสาขาวิชาที่เหมาะสมได้แม่นยำกว่า Naive Bayes เมื่อพิจารณาจากค่าความถูกต้อง (Accuracy)

งานวิจัยการพัฒนา โมเดลการแนะนำแนวทางการศึกษาของ [11] ใช้เทคนิคการทำเหมืองข้อมูล เช่น Decision Tree, Neural Networks และ Support Vector Machine เพื่อพัฒนาโมเดลแนะนำแขนงวิชา พบว่า Decision Tree มีความสามารถสูงในการแนะนำสาขาวิชาที่เหมาะสมกับนักศึกษา โดยมีความแม่นยำสูงสุดในกลุ่มวิศวกรรมซอฟต์แวร์ถึง 86.67% และงานวิจัยของ [12] วิเคราะห์ปัจจัยที่เกี่ยวข้องกับการเลือกศึกษาต่อในระดับมหาวิทยาลัย โดยใช้เทคนิค FP-

Growth และ Evolutionary Selection พบว่าปัจจัยสำคัญ เช่น ชื่อหลักสูตรและเนื้อหาของหลักสูตร มีผลต่อการตัดสินใจเลือกสาขาวิชาและผลสัมฤทธิ์ทางการเรียนของนักศึกษา นอกจากนี้งานวิจัยของ [13] ศึกษาความสัมพันธ์ระหว่างข้อมูลผลการเรียนระดับโรงเรียนเดิมและมหาวิทยาลัย พบว่า เกรดเฉลี่ยจากโรงเรียนเดิมในวิชาหลัก เช่น คณิตศาสตร์และวิทยาศาสตร์ สามารถใช้ร่วมกับข้อมูลในมหาวิทยาลัย เช่น การเข้าร่วมกิจกรรม เพื่อพัฒนาโมเดลแนะนำแนวการศึกษาได้อย่างมีประสิทธิภาพ

จากการศึกษางานวิจัยที่กล่าวมาข้างต้น พบว่าเทคนิคการทำเหมืองข้อมูลสามารถนำมาประยุกต์ใช้ในการทำนายสาขาวิชาที่เหมาะสมกับนักเรียนได้ ผู้วิจัยจึงมีแนวคิดในการใช้เทคนิคเหมืองข้อมูล ได้แก่ Decision Tree, Naïve Bayes, Support Vector Machine (SVM), Random Forest และ AdaBoost ในการพัฒนาโมเดลสำหรับทำนายสาขาวิชาที่เหมาะสมกับผู้สมัครเรียนที่มหาวิทยาลัยราชภัฏกำแพงเพชร โดยแตกต่างจากงานวิจัยก่อนหน้านี้ เนื่องจากผู้วิจัยใช้ข้อมูลจากมหาวิทยาลัยราชภัฏกำแพงเพชรและเพิ่มการวิเคราะห์ข้อมูลด้วย UMAP (Uniform Manifold Approximation and Projection) Principal Component Analysis (PCA) เพื่ออธิบายการกระจายตัวของข้อมูลในแต่ละระดับผลสัมฤทธิ์ โดยงานวิจัยนี้มีวัตถุประสงค์เพื่อ 1) พัฒนาโมเดลการทำนายสาขาวิชาที่เหมาะสมโดยใช้ Decision Tree, Naïve Bayes, SVM, Random Forest และ AdaBoost บนข้อมูลผลการเรียนจากโรงเรียนเดิมและในมหาวิทยาลัย 2 ชั้นปี 2) เปรียบเทียบประสิทธิภาพของโมเดลจำแนกประเภททั้ง 5 เทคนิค และ 3) วิเคราะห์ความสัมพันธ์ระหว่างผลการเรียนในโรงเรียนเดิมกับความสำเร็จในมหาวิทยาลัย

มหาวิทยาลัยราชภัฏกำแพงเพชร เป็นมหาวิทยาลัยกลุ่มที่ 3 ที่มุ่งเน้นการพัฒนาชุมชนท้องถิ่น

และชุมชนอื่น ๆ โดยมีปรัชญาการศึกษาที่ให้ความสำคัญกับการพัฒนาผู้เรียนตามผลลัพธ์การเรียนรู้ บูรณาการการเรียนการสอนกับการปฏิบัติจริง และใช้ชุมชนเป็นแหล่งเรียนรู้ผ่านกระบวนการทักษะวิศวกรสังคม ดังนั้นงานวิจัยนี้จึงมีความสำคัญต่อมหาวิทยาลัยราชภัฏกำแพงเพชร เนื่องจากช่วยให้นักศึกษาเลือกสาขาวิชาที่เหมาะสมกับตนเองมากขึ้น ซึ่งสอดคล้องกับแนวทางของมหาวิทยาลัยในการผลิตบัณฑิตที่มี ทักษะการเรียนรู้ตลอดชีวิต มีจิตอาสา และสามารถนำความรู้ไปพัฒนาท้องถิ่นได้อย่างมีประสิทธิภาพ

2. วัตถุประสงค์และวิธีการวิจัย

2.1 ข้อมูลที่ใช้ในการทดลอง

ในการวิจัยนี้ได้ทำการเก็บรวบรวมข้อมูลจากสำนักส่งเสริมวิชาการและงานทะเบียน มหาวิทยาลัยราชภัฏกำแพงเพชร โดยใช้ข้อมูลนักศึกษาระดับปริญญาตรีภาคปกติชั้นปีที่ 1 และชั้นปีที่ 2 ที่เข้าศึกษาปีการศึกษา 2559-2562 ทั้งหมด 10 หลักสูตร จำนวน 1,392 คน ซึ่งมีคลาสคำตอบจำนวน 4 คลาส ได้แก่ คลาสที่ 1 ผลการเรียนดีมาก (A) คะแนนเฉลี่ย 3.50-4.00 คลาสที่ 2 ผลการเรียนดี (B) คะแนนเฉลี่ยระหว่าง 3.00-3.49 คลาสที่ 3 ผลการเรียนปานกลาง (C) คะแนนเฉลี่ยระหว่าง 2.00-2.99 และคลาสที่ 4 ผลการเรียนต่ำ (E) คะแนนเฉลี่ยต่ำกว่า 2.00

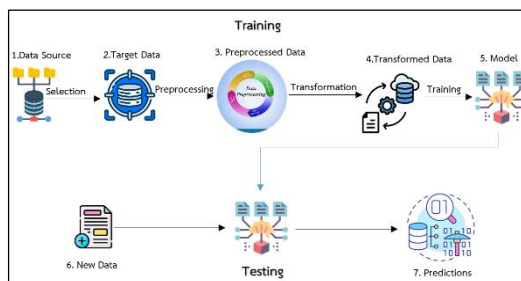
2.2 กระบวนการทำเหมืองข้อมูล

การออกแบบการทดลองเพื่อสร้างตัวแบบสำหรับการจำแนกคลาส มีการเปรียบเทียบเทคนิคสำหรับการสร้างตัวแบบในการจำแนกคลาส 5 เทคนิค โดยแบ่งข้อมูลแต่ละคลาสที่ใช้ในการทดลองออกเป็น 2 ชุด ได้แก่

Training Data (70%): ใช้สำหรับการสร้างโมเดล โดยโมเดลจะเรียนรู้จากข้อมูลในกลุ่มนี้เพื่อค้นหาความสัมพันธ์ระหว่างตัวแปรต้นและตัวแปรตาม

Testing Data (30%): ใช้สำหรับการประเมินความถูกต้องของโมเดล โดยข้อมูลในกลุ่มนี้ไม่ถูกใช้ในขั้นตอน

การสร้างโมเดลเพื่อให้การตรวจสอบมีความน่าเชื่อถือและสะท้อนความสามารถของโมเดลได้จริง [9] โดยมีขั้นตอนการทำวิจัยดังรูปที่ 1



รูปที่ 1 กระบวนการทดลอง

รูปที่ 1 มีขั้นตอนในการทดลองดังนี้

1) การรวบรวมข้อมูล ในงานวิจัยข้อมูลที่ใช้ในการทดลองได้มาจากสำนักส่งเสริมวิชาการและงานทะเบียน มีทั้งข้อมูลจากโรงเรียนเดิม และข้อมูลผลการเรียนในชั้นปีที่ 1 และปีที่ 2 ของนักศึกษาระดับปริญญาตรีภาคปกติ ที่เข้าศึกษาปีการศึกษา 2559-2562

2) การเลือกหลักสูตรเป้าหมายเพื่อนำมาใช้ในการทดลอง ผู้วิจัยเลือก 10 หลักสูตร มีรายละเอียดดังนี้ กลุ่มหลักสูตรวิทยาศาสตร์และเทคโนโลยี ได้แก่ หลักสูตรฟิสิกส์ หลักสูตรเคมี หลักสูตรชีววิทยา หลักสูตรคอมพิวเตอร์ และหลักสูตรสาธารณสุขศาสตร์ กลุ่มหลักสูตรมนุษยศาสตร์และสังคมศาสตร์ ได้แก่ หลักสูตรดนตรีศึกษา หลักสูตรการจัดการทั่วไป หลักสูตรนิติศาสตร์ หลักสูตรการศึกษาปฐมวัย และหลักสูตรบรรณารักษศาสตร์ เนื่องจากข้อมูลมีความสมบูรณ์ที่สุด และการกระจายตัวของคลาสเป้าหมายสูง

3) การเตรียมข้อมูล ข้อมูลที่ได้จากขั้นตอนที่ 2 ผู้วิจัยทำความสะอาดข้อมูล ตัดข้อมูลที่ไม่มีสมบูรณ์ออกเหลือข้อมูลสำหรับทดลองจำนวน 1,392 รายการ โดยแต่ละรายการประกอบด้วยแอทริบิวต์ ดังนี้ เกรดจากโรงเรียน

เดิม (ม.4-ม.6) ในรายวิชาคณิตศาสตร์, วิทยาศาสตร์, ภาษาอังกฤษ, ภาษาไทย, สังคม, ศิลปะ จากนั้น แบ่งข้อมูลในแต่ละคลาสออกเป็น 2 กลุ่มคือข้อมูลสอนระบบและทดสอบระบบ ด้วยอัตราส่วน 70 : 30

4) การเปลี่ยนรูปแบบข้อมูลให้อยู่ในรูปแบบที่โมเดลเหมืองข้อมูลเข้าใจผู้วิจัยใช้เทคนิค Label Encoding เช่น จังหวัดกำแพงเพชร ตาก สุโขทัย แทนที่ด้วย 2 11 และ 45 เป็นต้น ซึ่งมีรายละเอียดคุณลักษณะข้อมูลดังแสดงในตารางที่ 1

ตารางที่ 1 คุณลักษณะแอตทริบิวต์ที่ใช้ในการทดลอง

ชื่อแอตทริบิวต์	รายละเอียด	ชนิดข้อมูล
Gender	เพศ (1 = ชาย, 2 = หญิง)	Numeric
Thai_GPA	เกรดเฉลี่ยรายวิชาภาษาไทยตั้งแต่ ม.4-ม.6	Numeric
Math_GPA	เกรดเฉลี่ยรายวิชาคณิตศาสตร์ตั้งแต่ ม.4-ม.6	Numeric
Sci_GPA	เกรดเฉลี่ยรายวิชาวิทยาศาสตร์ตั้งแต่ ม.4-ม.6	Numeric
Social_GPA	เกรดเฉลี่ยรายวิชาสังคมศึกษาตั้งแต่ ม.4-ม.6	Numeric
Health_GPA	เกรดเฉลี่ยรายวิชาสุขภาพตั้งแต่ ม.4-ม.6	Numeric
Art_GPA	เกรดเฉลี่ยรายวิชาศิลปะตั้งแต่ ม.4-ม.6	Numeric
Career_GPA	เกรดเฉลี่ยรายวิชาการงานอาชีพตั้งแต่ ม.4-ม.6	Numeric
Eng_GPA	เกรดเฉลี่ยรายวิชาภาษาอังกฤษตั้งแต่ ม.4-ม.6	Numeric
preSchool	โรงเรียนเดิมที่ศึกษาในระดับ ม.4-ม.6	Numeric
provinceSchool	จังหวัดของโรงเรียนเดิมที่ศึกษาในระดับ ม.4-ม.6	Numeric
Faculty	คณะ	Numeric
Major	หลักสูตร	Numeric
class	GPA จนถึงชั้นปีที่ 2 (คลาสผลลัพธ์)	Categorical

ตารางที่ 1 แสดงคุณลักษณะแอตทริบิวต์ที่ใช้ในการทดลอง แบ่งเป็นแอตทริบิวต์สำหรับสอนระบบได้แก่แอตทริบิวต์ Gender ถึง Major และแอตทริบิวต์ class เป็นเป้าหมายการเรียนรู้ซึ่งมีทั้งหมด 4 คลาส ได้แก่ A, B, C และ E ส่วนแอตทริบิวต์ GPAของแต่ละรายวิชาผู้วิจัยได้ทำการแปลงรูปแบบจากเลขทศนิยมเป็นตัวเลขจำนวนเต็มโดยกำหนดเงื่อนไขการแปลงรูปแบบดังนี้

เกรดเฉลี่ย ≤ 0.99 แปลงเป็นเลข 1

5) การสร้างโมเดล เทคนิคที่ใช้คือ Decision Tree, Naive Bayes, SVM, Random Forest และ AdaBoost เพื่อเปรียบเทียบประสิทธิภาพ ในการจำแนกประเภทคลาสผลลัพธ์ที่ได้คือโมเดลที่นำมาใช้กับข้อมูลทดสอบ

6) การทดสอบประสิทธิภาพของโมเดลด้วยข้อมูลทดสอบระบบ โดยใช้วิธี cross-validation แบบวนรอบทั้งหมด 10 ครั้ง

เกรดเฉลี่ย 1.00-1.49 แปลงเป็นเลข 2

เกรดเฉลี่ย 1.50-1.99 แปลงเป็นเลข 3

เกรดเฉลี่ย 2.00-2.49 แปลงเป็นเลข 4

เกรดเฉลี่ย 2.50-2.99 แปลงเป็นเลข 5

เกรดเฉลี่ย 3.00-3.49 แปลงเป็นเลข 6

เกรดเฉลี่ย 3.50-3.99 แปลงเป็นเลข 7

เกรดเฉลี่ย 4.00 แปลงเป็นเลข 8

2.3 เครื่องมือที่ใช้ในการทดลอง

เครื่องคอมพิวเตอร์หน่วยประมวลผลกลาง Core i7 หน่วยความจำหลัก 16 GB ซอฟต์แวร์ด้านการทำเหมืองข้อมูล WEKA 3.8.6

2.4 ตัววัดในการประเมินโมเดล

การประเมินโมเดลในงานวิจัยนี้ใช้ตัววัดที่หลากหลายเพื่อวิเคราะห์ประสิทธิภาพในมิติต่าง ๆ [14] ได้แก่

1) Precision (ความถูกต้อง): วัดความถูกต้องของผลลัพธ์ในกลุ่มที่โมเดลทำนายว่าเป็น Positive

$$Precision = \frac{TP}{TP+FP} \quad (1)$$

2) Recall (ความครอบคลุม): วัดความสามารถของโมเดลในการดึงข้อมูลที่เป็น Positive จริงทั้งหมด

$$Recall = \frac{TP}{TP+FN} \quad (2)$$

3) F1-Score: คำนวณค่าเฉลี่ยเชิงฮาร์โมนิกระหว่าง Precision และ Recall

$$F1\ Score = \frac{2 \times Precision \times Recall}{Precision+Recall} \quad (3)$$

โดยที่ TP คือ ค่าที่ทำนายตรงกับค่าจริง และถูกต้อง

FP คือ ค่าที่ทำนายไม่ตรงกับค่าจริง

FN คือ ค่าที่ทำนายไม่ตรงกับค่าจริง และไม่ถูกต้อง

3. ผลการทดลองและอภิปรายผล

ตารางที่ 2 และ ตารางที่ 3 แสดงผลผลการประเมินประสิทธิภาพของโมเดลในขั้นตอนการฝึกโมเดล ซึ่งมีโมเดลที่ใช้ในการทดลองได้แก่ Decision Tree (M1), Naive Bayes (M2), SVM (M3), Random Forest (M4) และ AdaBoost (M5) โดยวัดประสิทธิภาพด้วยความแม่นยำในการเรียนรู้ของโมเดล (Train Accuracy) ตารางที่ 3 ผลประเมินประสิทธิภาพของโมเดลในขั้นตอนการทดสอบกลุ่มหลักสูตรตัวอย่าง 10 หลักสูตรประเมินประสิทธิภาพด้วยการวัดค่า Precision, Recall และ F1-Score

ตารางที่ 2 ผลประเมินประสิทธิภาพของโมเดลในขั้นตอนการฝึกโมเดล

	Train Accuracy				
	M1	M2	M3	M4	M5
ฟิสิกส์	60.80	64.00	64.00	61.60	59.2
เคมี	54.47	59.39	52.03	58.54	56.91
ชีววิทยา	35.48	29.03	29.03	22.58	35.48
คอมพิวเตอร์	51.01	54.25	57.49	57.08	52.63
สาธารณสุขศาสตร์	53.85	59.34	60.44	59.34	46.15
ดนตรีศึกษา	55.56	55.56	41.27	61.90	52.38
การจัดการทั่วไป	56.25	62.50	64.58	60.42	62.5
นิเทศศาสตร์	70.67	69.33	78.67	74.67	78.67
การศึกษารัฐบาล	48.04	50.00	47.06	50.98	39.22
บรรณารักษศาสตร์	66.67	60.87	72.46	66.67	60.87

ตารางที่ 3 ผลประเมินประสิทธิภาพของโมเดลในขั้นตอนการทดสอบกลุ่มหลักสูตรตัวอย่าง 10 หลักสูตร

	Test														
	Precision					Recall					F1-Score				
	M1	M2	M3	M4	M5	M1	M2	M3	M4	M5	M1	M2	M3	M4	M5
ฟิสิกส์	0.61	0.67	0.54	0.67	0.35	0.61	0.65	0.61	0.67	0.59	0.61	0.65	0.55	0.65	0.44
เคมี	0.52	0.56	0.55	0.52	0.46	0.50	0.60	0.50	0.56	0.50	0.50	0.56	0.52	0.53	0.41
ชีววิทยา	0.28	0.52	0.37	0.65	0.24	0.31	0.54	0.46	0.62	0.46	0.29	0.47	0.38	0.63	0.32
คอมพิวเตอร์	0.47	0.54	0.39	0.51	0.39	0.44	0.54	0.51	0.51	0.49	0.45	0.53	0.43	0.51	0.43
สาขารณสุขศาสตร์	0.60	0.58	0.63	0.61	0.34	0.54	0.59	0.62	0.49	0.41	0.54	0.53	0.59	0.52	0.37
ดนตรีศึกษา	0.44	0.59	0.58	0.43	0.60	0.48	0.68	0.60	0.48	0.68	0.45	0.63	0.57	0.45	0.61
การจัดการทั่วไป	0.50	0.63	0.38	0.71	0.38	0.57	0.71	0.62	0.81	0.62	0.53	0.66	0.47	0.75	0.47
นิเทศศาสตร์	0.61	0.65	0.61	0.60	0.61	0.78	0.72	0.78	0.72	0.78	0.69	0.68	0.69	0.65	0.69
การศึกษารัฐมวัย	0.53	0.58	0.47	0.61	0.40	0.50	0.66	0.32	0.59	0.45	0.51	0.61	0.25	0.58	0.42
บรรณารักษศาสตร์	0.76	0.73	0.51	0.67	0.73	0.70	0.73	0.47	0.67	0.70	0.71	0.73	0.46	0.67	0.71

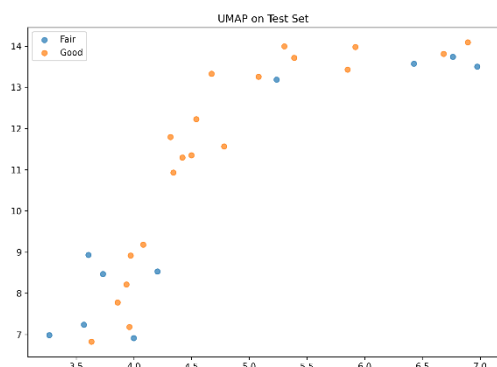
ตารางที่ 2 แสดงค่า ความแม่นยำในการฝึก (Train Accuracy) ของโมเดลการจำแนกประเภทจำนวน 5 โมเดล ได้แก่ Decision Tree (M1), Naive Bayes (M2), Support Vector Machine (M3), Random Forest (M4) และ AdaBoost (M5) โดยทำการประเมินบนข้อมูลจาก 10 หลักสูตร จากผลการประเมิน พบว่า Naive Bayes (M2) มีค่า Train Accuracy สูงสุดในบางหลักสูตร ได้แก่ ฟิสิกส์ (64.00) และเคมี (59.39) ซึ่งสอดคล้องกับสมมติฐานพื้นฐานของโมเดล Naive Bayes ที่เหมาะกับข้อมูลที่ตัวแปรมีความเป็นอิสระต่อกัน (Independent Features) และมีโครงสร้างของกลุ่มข้อมูลที่ไม่ซับซ้อน ในขณะที่ SVM (M3) แสดงค่า Train Accuracy สูงสุดในหลายหลักสูตรที่มีความซับซ้อนของข้อมูลมากขึ้น เช่น คอมพิวเตอร์ (57.49) สาขารณสุขศาสตร์ (60.44) การจัดการทั่วไป (64.58) และ นิเทศศาสตร์ (78.67) ซึ่งแสดงถึงศักยภาพของ SVM ในการเรียนรู้จากข้อมูลที่มีเส้นแบ่งเชิงซ้อนหรือมีมิติสูง สำหรับโมเดลในกลุ่ม Ensemble ได้แก่ Random Forest (M4) และ AdaBoost (M5) พบว่าให้ผลการฝึกที่มีเสถียรภาพในระดับปานกลางถึงสูงในหลายหลักสูตร

โดยเฉพาะในกรณีที่มีข้อมูลมีลักษณะแปรผันหรือมีความหลากหลายภายในกลุ่ม เช่น ชีววิทยา และ การศึกษาปฐมวัย

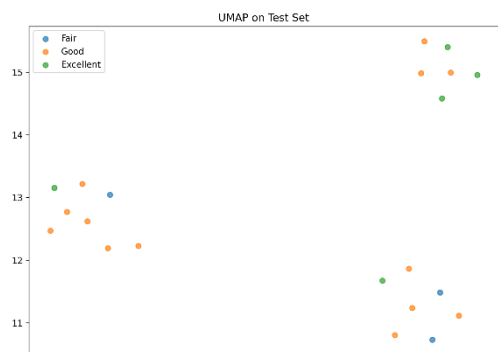
จากผลการทดลองที่แสดงใน ตารางที่ 3 พบว่าโมเดล Naive Bayes (M2) แสดงประสิทธิภาพโดยรวมได้ดีที่สุด โดยมีค่า F1-Score สูงสุดใน 6 จาก 10 หลักสูตร ซึ่งสะท้อนถึงความเหมาะสมของโมเดลกับข้อมูลที่มีโครงสร้างไม่ซับซ้อน หรือมีความแตกต่างระหว่างกลุ่มเป้าหมายอย่างชัดเจน เช่น ในหลักสูตร ดนตรีศึกษา การศึกษารัฐมวัย และบรรณารักษศาสตร์ ทั้งนี้ แม้ Naive Bayes จะให้ผลลัพธ์ที่โดดเด่นโดยรวม แต่โมเดลอื่น เช่น Random Forest และ AdaBoost ยังแสดงศักยภาพในการจำแนกได้ดีในบางบริบท โดยเฉพาะในหลักสูตรที่มีข้อมูลมีลักษณะซับซ้อนหรือมีความแปรผันสูง เช่น การจัดการทั่วไป และ นิเทศศาสตร์ ซึ่งอาจสะท้อนถึงประสิทธิภาพของโมเดลในบริบทที่ต้องจัดการกับข้อมูลที่มีลักษณะไม่สมดุลหรือมีขอบเขตการจำแนกไม่ชัดเจน [15]

นอกจากนี้ผู้วิจัยดำเนินการวิเคราะห์เพิ่มเติมโดยใช้เทคนิคการลดมิติ (Dimensionality Reduction) ซึ่งเป็น

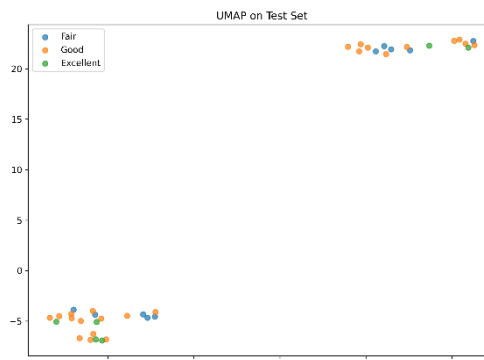
กระบวนการแปลงข้อมูลจากคุณลักษณะหลายมิติให้อยู่ในรูปแบบ สองมิติ เพื่อให้สามารถสังเกตโครงสร้าง การกระจายตัว และรูปแบบของข้อมูลได้อย่างชัดเจนยิ่งขึ้น โดยในงานวิจัยนี้เลือกใช้เทคนิค UMAP (Uniform Manifold Approximation and Projection) ซึ่งเป็นเทคนิคลดมิติแบบไม่เป็นเชิงเส้นที่มีประสิทธิภาพสูง และได้รับความนิยมในงานด้านการวิเคราะห์ข้อมูลและการเรียนรู้ของเครื่อง เนื่องจากสามารถรักษาโครงสร้างของข้อมูลทั้งในระดับภาพรวม (global structure) และระดับกลุ่มย่อย (local structure) ได้อย่างมีประสิทธิภาพ [16] ด้วยเหตุนี้ UMAP จึงเหมาะสมอย่างยิ่งสำหรับการสำรวจและตีความข้อมูลเชิงจำแนกในแต่ละบริบทของหลักสูตรดังกล่าวอย่างในรูปที่ 2 – รูปที่ 4



รูปที่ 2 การกระจายของข้อมูลทดสอบตามระดับผลสัมฤทธิ์ทางการเรียนหลักสูตรบรรณารักษศาสตร์ด้วยเทคนิค UMAP



รูปที่ 3 การกระจายของข้อมูลทดสอบตามระดับผลสัมฤทธิ์ทางการเรียนหลักสูตรการจัดการทั่วไปด้วยเทคนิค UMAP



รูปที่ 4 การกระจายของข้อมูลทดสอบตามระดับผลสัมฤทธิ์ทางการเรียนหลักสูตรสาธารณสุขศาสตร์ด้วยเทคนิค UMAP

การวิเคราะห์การกระจายของข้อมูลชุดทดสอบด้วยเทคนิค UMAP ในรูปที่ 2, 3 และ 4 แสดงให้เห็นความแตกต่างเชิงโครงสร้างของข้อมูลในแต่ละหลักสูตร ได้แก่ บรรณารักษศาสตร์, การจัดการทั่วไป และ สาธารณสุขศาสตร์ ตามลำดับ ซึ่งสามารถใช้ประกอบการอภิปรายถึงสาเหตุที่โมเดลบางประเภทสามารถทำนายผลสัมฤทธิ์ทางการเรียนได้อย่างมีประสิทธิภาพ ใน รูปที่ 2 ซึ่งแสดงข้อมูลของหลักสูตรบรรณารักษศาสตร์ พบว่าการกระจายของข้อมูลแบ่งออกเป็นสองกลุ่มหลักตามแนวโค้ง โดยมีการซ้อนทับกันบางส่วนระหว่างระดับ Fair และ Good แต่ยังคงรักษาโครงสร้างการแยกกลุ่มได้พอสมควร ลักษณะดังกล่าวเหมาะสมกับโมเดลที่สามารถจัดการกับข้อมูลที่มีโครงสร้างง่ายและรองรับการจำแนกเชิงน่าจะเป็น เช่น Naive Bayes (M2) ซึ่งจากผลใน ตารางที่ 3 พบว่าโมเดลนี้ให้ค่า F1-Score สูงที่สุดในหลักสูตรนี้ (0.73) ซึ่งสอดคล้องกับข้อเสนอของ Hastie et al. [15] ที่ระบุว่า Naive Bayes เหมาะกับข้อมูลที่มีการแยกกลุ่มชัดเจนและมีลักษณะการแจกแจงที่เป็นอิสระ

ใน รูปที่ 3 ซึ่งเป็นข้อมูลของหลักสูตรการจัดการทั่วไป พบว่ามีการกระจายตัวที่แบ่งกลุ่มได้อย่างเด่นชัด โดยเฉพาะกลุ่ม Excellent ที่อยู่ห่างจากกลุ่มอื่นอย่างชัดเจน ขณะที่กลุ่ม Good และ Fair มีแนวโน้มกระจายแยกกันในเชิงแนวราบ ลักษณะดังกล่าวสะท้อนว่า

โครงสร้างข้อมูลมีการแยกคลาสที่ชัดเจน ซึ่งเอื้อต่อการเรียนรู้ของโมเดลที่มีความสามารถในการจับโครงสร้างแบบไม่เชิงเส้นและจัดการกับความแปรปรวนในข้อมูลได้ดี เช่น Random Forest (M4) ที่ให้ค่า F1-Score สูงสุดในหลักสูตรนี้ (0.75) ซึ่งสนับสนุนคุณสมบัติของโมเดล Ensemble ในการจัดการกับข้อมูลที่ซับซ้อน [15]

รูปที่ 4 ที่แสดงข้อมูลของหลักสูตรสาธารณสุขศาสตร์ พบว่าข้อมูลกระจายออกเป็นสองกลุ่มใหญ่ที่มีระยะห่างชัดเจน แต่ภายในกลุ่มยังพบการซ้อนทับระหว่างคลาส Good, Fair และ Excellent ลักษณะการแยกกลุ่มที่ไม่เป็นเชิงเส้นและมีขอบเขตคลุมเครือเช่นนี้ ทำให้โมเดลที่สามารถสร้างเส้นขอบเขตที่ยืดหยุ่น เช่น Support Vector Machine (M3) ทำงานได้ดีขึ้น โดยจากผลใน ตารางที่ 3 พบว่า SVM ให้ค่า F1-Score สูงสุดในหลักสูตรนี้ (0.59) ซึ่งสอดคล้องกับแนวคิดของการใช้ Kernel-Based Classifier ในการจำแนกข้อมูลที่ไม่สามารถแยกเส้นตรงได้ในมิติเดิม [17] โดยสรุป ความแตกต่างของโครงสร้างข้อมูลที่ปรากฏจากการวิเคราะห์ด้วย UMAP ช่วยอธิบายถึงความเหมาะสมของโมเดลจำแนกประเภทในแต่ละบริบทของหลักสูตรได้อย่างชัดเจน และสะท้อนว่าการเลือกโมเดลไม่ควรพิจารณาจากค่าเฉลี่ยโดยรวมเท่านั้น แต่ต้องอิงกับลักษณะเชิงโครงสร้างของข้อมูลแต่ละกลุ่มด้วย [16]

4. สรุปผล

การศึกษานี้มุ่งพัฒนาโมเดลเพื่อทำนายความเหมาะสมของสาขาวิชาในระดับปริญญาตรี โดยประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่อง 5 ประเภท ได้แก่ Decision Tree (M1), Naive Bayes (M2), Support Vector Machine (SVM, M3), Random Forest (M4) และ AdaBoost (M5) บนชุดข้อมูลนักศึกษาจาก 10 หลักสูตรในมหาวิทยาลัยราชภัฏกำแพงเพชร ผลลัพธ์ที่ได้จากโมเดลจะทำนายว่านักศึกษาที่เลือกเรียนในสาขาใดจะสามารถเรียนจนสำเร็จการศึกษาได้หรือไม่และ

จะเรียนสำเร็จด้วยระดับใด การทำนายผลสัมฤทธิ์ทางการเรียนจะให้ผลลัพธ์เป็น 4 กลุ่มระดับผลสัมฤทธิ์ ได้แก่ ระดับดีมาก (Excellent) ระดับดี (Good) ระดับปานกลาง (Fair) และ ไม่ผ่าน (Fail) การวิเคราะห์ผลการทดลองพบว่าโมเดล Naive Bayes (M2) ให้ผลลัพธ์โดยรวมดีที่สุด โดยได้ค่า F1-Score สูงสุดใน 6 จาก 10 หลักสูตร ซึ่งสอดคล้องกับลักษณะข้อมูลที่มีการแยกกลุ่มระหว่างคลาสได้ชัดเจน และคุณลักษณะต่าง ๆ มีแนวโน้มเป็นอิสระต่อกัน (feature independence) เช่น ในหลักสูตร ดนตรีศึกษา, การศึกษาปฐมวัย, และบรรณารักษศาสตร์ โมเดล SVM (M3) แสดงศักยภาพสูงในหลักสูตรที่มีโครงสร้างข้อมูลซับซ้อน เช่น สาธารณสุขศาสตร์ ซึ่งข้อมูลมีแนวโน้มซ้อนทับระหว่างกลุ่มผลสัมฤทธิ์ การใช้ kernel function ช่วย SVM สามารถสร้างเส้นขอบเขตการจำแนกที่ยืดหยุ่นในมิติสูงได้ดี ด้าน Random Forest (M4) และ AdaBoost (M5) ซึ่งเป็นโมเดลกลุ่ม Ensemble Learning ก็แสดงประสิทธิภาพเด่นในหลักสูตรที่มีการแปรผันของข้อมูลสูง เช่น การจัดการทั่วไป และ นิติศาสตร์ นอกจากนี้ การวิเคราะห์เชิงโครงสร้างด้วยเทคนิค UMAP (Uniform Manifold Approximation and Projection) ช่วยให้เห็นภาพการกระจายของข้อมูลในแต่ละระดับผลสัมฤทธิ์ได้อย่างชัดเจน

แต่ถึงอย่างไรผลจากการทดลองดัง ตารางที่ 2 และ ตารางที่ 3 ชี้ให้เห็นว่าข้อมูลพื้นฐานและผลการเรียนจากโรงเรียนเดิมไม่ได้ส่งผลโดยตรงต่อความสำเร็จในการเรียนที่มหาวิทยาลัยราชภัฏกำแพงเพชร ซึ่งอาจจำเป็นต้องใช้ข้อมูลเพิ่มเติมเพื่อการทำนายมีความแม่นยำยิ่งขึ้น สอดคล้องกับงานวิจัยของ [8] [18] ที่เสนอให้มีการจัดเก็บประเภทและรายละเอียดของข้อมูลเพิ่มเติม จากการศึกษาของผู้วิจัยพบว่าปัจจัยอื่น ๆ เช่น พฤติกรรมการเรียนรู้ เช่น ความสม่ำเสมอในการเข้าเรียน การทำแบบฝึกหัด การมีส่วนร่วมในกิจกรรมการเรียนการสอน รวมถึงการวัดผล

ทางการเรียนที่ไม่ใช่เพียงแค่เกรดเฉลี่ยหรือผลการทดสอบ อาจช่วยให้โมเดลมีข้อมูลที่ครอบคลุมมากขึ้น ในการทำนายและแนะนำสาขาวิชาที่เหมาะสมกับ ผู้สมัครเรียนแต่ละคนนอกจากนี้ การนำโมเดลไปใช้จริงในระบบแนะนำแนวของมหาวิทยาลัยจะช่วยให้ นักศึกษาปีแรกได้รับคำแนะนำที่เหมาะสมในการเลือก สาขาวิชา ซึ่งจะช่วยลดอัตราการออกกลางคัน สนับสนุนการตัดสินใจของนักศึกษาให้เลือกสาขาที่ ตรงกับความถนัด และเพิ่มโอกาสในการประสบความสำเร็จในการเรียนต่อไป

5. กิตติกรรมประกาศ

ขอขอบคุณมหาวิทยาลัยราชภัฏกำแพงเพชรในการสนับสนุนงบประมาณสำหรับการทำวิจัย

ขอขอบคุณสำนักส่งเสริมวิชาการและงานทะเบียน มหาวิทยาลัยราชภัฏกำแพงเพชร ที่ให้ความอนุเคราะห์ในการให้ข้อมูล

6. เอกสารอ้างอิง

- [1] Academic Promotion and Registration Office, Report: Student Academic Performance Data 2020 of Kamphaeng Phet Rajabhat University, 2020 (in Thai).
- [2] W. Matchai and W. Buathong, "A study of factors affecting the selection of vocational certificate students to continue in higher vocational education using data mining techniques, Chandrakasem Rajabhat University Journal of Science and Academic Research, vol. 33, no. 1, pp. 1-10, Jan.-Jun. 2023 (in Thai).

- [3] C. Wilkin, A. Ferreira, K. Rotaru, and L. R. Gaerlan, "Big data prioritization in SCM decision-making: Its role and performance implications," *International Journal of Accounting Information Systems*, vol. 100470, 2020. [Online]. Available: <https://doi.org/10.1016/j.accinf.2020.100470>.
- [4] I. A. T. Hashem, I. Yaqoob, N. B. Anuar, S. Mokhtar, A. Gani, and S. U. Khan, "The rise of 'big data' on cloud computing: Review and open research issues," Report, *Information Systems*, vol. 47, pp. 98–115, 2015. [Online]. Available: <https://doi.org/10.1016/j.is.2014.07.006>.
- [5] S. Ahmed, M. A. Khan, and M. Rashid, " Classification techniques in data mining for student performance prediction," *Computers in Human Behavior*, vol. 85, pp. 78–85, 2018.
- [6] C. Romero and S. Ventura, "Educational data mining: A review of the state of the art," *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, vol. 50, no. 3, pp. 601–618, 2020.
- [7] E. Osmanbegovic and M. Suljic, "Data mining approach for predicting student performance," *Economic Review – Journal of Economics and Business*, vol. 10, no. 1, pp. 3–12, May 2012.
- [8] H. Yao, D. Lian, Y. Cao, Y. Wu, and T. Zhou, "Predicting academic performance for college students: A campus behavior perspective," *ACM Transactions on Intelligent Systems and Technology*, vol. 10, pp. 1–21, 2019.
- [9] J. R. Quinlan, " Induction of Decision Trees," *Machine Learning*, vol. 1, no. 1, pp. 81–106, 1986.

- [10] A. Pinathe, "The Use of Data Mining in Selecting Areas of Study for Further Education Opportunity," *Journal of Science and Technology Mahasarakham University*, vol. 36, no. 6, pp. 704–712, 2017 (in Thai).
- [11] J. Laohawaranan, R. Limsuthiwanphoom, and B. Thanasophon, " Using data mining techniques for classifying and selecting study branches for technology faculty students," *KMITL Journal of Information Technology*, vol. 4, no. 2, pp. 1–9, 2015 (in Thai).
- [12] N. Napaporn Chongkasikit, " An Application of FP- Growth Algorithm to Find Factors in Choosing to Study in the Faculty of Industrial Technology Lampang Rajabhat University," *Industrial Technology Lampang Rajabhat University Journal*, vol. 11, no. 2, pp. 29–39, 2018 (in Thai).
- [13] K. Hengpraprom, S. Hengpraprom, and S. Makwiboonchai, " The Knowledge Discovery of Students' Critical Characteristics towards Learning Achievement of Nakhon Pathom Rajabhat University Students' Program in Computer via Data Mining Technique," *Journal of Western Rajabhat Universitie*, vol. 9, no. 1, pp. 71–80, 2014 (in Thai).
- [14] P. Eakasit, *An Introduction to Data Mining Techniques*, Cube Edu-Scene, 2015 (in Thai).
- [15] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, 2nd ed. New York: Springer, 2009.
- [16] L. McInnes and J. Healy, "UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction," *arXiv preprint arXiv:1802.03426*, 2018. [Online]. Available: <https://doi.org/10.48550/arXiv.1802.03426>
- [17] V. Vapnik, *Statistical Learning Theory*, New York: Wiley, 1998.
- [18] P. Phetkong and C. Biaokaimuk, "MACHINE LEARNING BASED PREDICTION ANALYSIS FOR THE SERVICE LEVEL AGREEMENT: A CASE STUDY OF DHANARAK ASSET DEVELOPMENT CO., LTD," *SAU Journal of Science & Technology*, vol. 9, no. 2, pp. 66–76, Jul.–Dec. 2023 (in Thai).

ประวัติผู้ประพันธ์ :

**กนกวรรณ เขียววัน**

อาจารย์ประจำหลักสูตรวิทยา
ศาสตรบัณฑิต สาขาวิชาเทคโนโลยี
คอมพิวเตอร์ คณะเทคโนโลยี

อุตสาหกรรม มหาวิทยาลัยราชภัฏกำแพงเพชร
วท.ม. เทคโนโลยีสารสนเทศ สถาบันเทคโนโลยี
พระจอมเกล้าเจ้าคุณทหารลาดกระบัง

**จัมมิษา ตันคั่นติสม**

อาจารย์ประจำหลักสูตรวิทยา
ศาสตรบัณฑิต สาขาวิชาเทคโนโลยี
สารสนเทศ คณะวิทยาศาสตร์

มหาวิทยาลัยราชภัฏกำแพงเพชร
DIT (Information Technology) Edith Cowan
University

**นรุตม์ บุตรพลอย**

อาจารย์ประจำหลักสูตรวิทยา
ศาสตรบัณฑิต สาขาวิชาเทคโนโลยี
คอมพิวเตอร์ คณะเทคโนโลยี

อุตสาหกรรม มหาวิทยาลัยราชภัฏกำแพงเพชร
ปร.ด. วิศวกรรมคอมพิวเตอร์ มหาวิทยาลัยขอนแก่น

**พรหมเมศ วีระพันธ์**

อาจารย์ประจำหลักสูตรวิทยา
ศาสตรบัณฑิต สาขาวิชาเทคโนโลยี
สารสนเทศ คณะวิทยาศาสตร์

มหาวิทยาลัยราชภัฏกำแพงเพชร
วท.ม. เทคโนโลยีสารสนเทศ มหาวิทยาลัยนเรศวร

**จินดาพร อ่อนเกตุ**

อาจารย์ประจำหลักสูตรวิทยา
ศาสตรบัณฑิต สาขาวิชาเทคโนโลยี
สารสนเทศ คณะวิทยาศาสตร์

มหาวิทยาลัยราชภัฏกำแพงเพชร
วท.ม. เทคโนโลยีสารสนเทศและการจัดการ
มหาวิทยาลัยเชียงใหม่

**คมกริช กลิ่นอาจ**

หัวหน้ากลุ่มงานเทคโนโลยี
สารสนเทศ สำนักส่งเสริมวิชาการ
และงานทะเบียน มหาวิทยาลัย

ราชภัฏกำแพงเพชร
วท.ม. เทคโนโลยีสารสนเทศ มหาวิทยาลัยนเรศวร