

การสร้างตัวแบบและเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการจำแนก ประเภทโรคมะเร็งลำไส้โดยใช้ภาพทางจุลพยาธิวิทยา Building a Model and Evaluating the Performance for Colon Cancer Screening Using Histopathological Images

เลอศักดิ์ โพธิ์ทอง ชรัญญา ภารประสาท ปฏิกาน ทองอยู่ และ อนุนพงศ์ สุขประเสริฐ*

Lersak Phothong Charanya Phanprasat Patiparn Thongyu and Anupong Sukprasert*

สาขาวิชาคอมพิวเตอร์ธุรกิจ คณะการบัญชีและการจัดการ มหาวิทยาลัยมหาสารคาม จ.มหาสารคาม ประเทศไทย

Department of Business Computer, Mahasarakham Business School, Mahasarakham University,

Mahasarakham, Thailand

วันที่ส่งบทความ : 3 ธันวาคม 2566 วันที่แก้ไขบทความ : 15 ธันวาคม 2568 วันที่ตอบรับบทความ : 17 ธันวาคม 2568

Received: 3 December 2023, Revised: 15 December 2025, Accepted: 17 December 2025

บทคัดย่อ

มะเร็งลำไส้ใหญ่ยังคงเป็นข้อกังวลด้านสุขภาพที่สำคัญทั่วโลก การตรวจพบตั้งแต่ระยะเริ่มต้นมีความสำคัญอย่างยิ่งต่อการปรับปรุงผลลัพธ์ของการรักษา กระบวนการวินิจฉัยที่ผ่านมาแม้จะมีคุณภาพเป็นที่ยอมรับแต่ก็อาจใช้ระยะเวลาค่อนข้างมากและขึ้นกับประสบการณ์ของผู้เชี่ยวชาญ การวิจัยครั้งนี้มีวัตถุประสงค์เพื่อทดสอบประสิทธิภาพของตัวแบบการเรียนรู้ของเครื่องในการจำแนกโรคมะเร็งลำไส้ ด้วยภาพทางจุลพยาธิวิทยาจำนวน 10,000 ภาพ ซึ่งเผยแพร่ในเว็บไซต์ www.kaggle.com จากนั้นนำมาวิเคราะห์ตามกระบวนการมาตรฐานการทำเหมืองข้อมูล โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพจำนวน 4 เทคนิค ได้แก่ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) เทคนิคเพื่อนบ้านใกล้ที่สุด (k-Nearest Neighbors: k-NN) เทคนิคโครงข่ายประสาทเทียม (Neural Network: NN) และเทคนิคต้นไม้ตัดสินใจ (Decision Tree: DT) ผลการวิจัยพบว่า เทคนิค k-NN ให้ค่าความแม่นยำ (Accuracy) สูงที่สุด มีค่าเท่ากับ 91.86% รวมถึงให้ค่าความไว (Sensitivity) และค่าความจำเพาะ (Specificity) สูงที่สุด คือ 92.24% และ 91.48% ตามลำดับ จึงบ่งชี้ถึงศักยภาพของเทคนิค k-NN ว่ามีความเหมาะสมสำหรับการนำมาสร้างตัวแบบสำหรับการจำแนกประเภทโรคมะเร็งลำไส้ ซึ่งนำไปสู่การพัฒนาเครื่องมือที่สำคัญในการตรวจหาโรคตั้งแต่ระยะเริ่มต้นและสนับสนุนการวางแผนการรักษาที่มีประสิทธิภาพมากขึ้น

คำสำคัญ : การจำแนกประเภทของภาพ การจำแนกเซลล์มะเร็งลำไส้ใหญ่ การเรียนรู้ของเครื่อง เทคนิคซัพพอร์ตเวกเตอร์แมชชีน เพื่อนบ้านใกล้ที่สุด

*ที่อยู่ติดต่อ Email address: pinkaew.s@tsu.ac.th

<https://doi.org/10.55003/scikmitl.2025.261425>

Abstract

Colon cancer remains a major global health concern. Early detection is critical for improving treatment outcomes. Although conventional diagnostic methods are generally reliable, they can be time-consuming and heavily dependent on the expertise of medical professionals. This study aims to evaluate the effectiveness of machine learning models in classifying colon cancer using 10,000 histopathological images obtained from www.kaggle.com. The data were analyzed following standard data mining procedures using four image classification techniques: Support Vector Machine (SVM), k-Nearest Neighbors (k-NN), Neural Network (NN), and Decision Tree (DT). The results showed that the k-NN technique achieved the highest accuracy at 91.86%, along with the highest sensitivity and specificity values at 92.24% and 91.48%, respectively. These findings indicate that the k-NN technique is highly suitable for developing classification models for colon cancer, contributing to the creation of essential tools for early detection and supporting more effective treatment planning.

Keywords: Image classification, Colon cancer classification, Machine Learning, Support Vector Machine, k-Nearest Neighbors

1. บทนำ

โรคมะเร็งลำไส้ใหญ่เป็นสาเหตุการเสียชีวิตอันดับที่สองจากสถิติการเสียชีวิตจากโรคมะเร็งทั่วโลก โดยข้อมูลใน พ.ศ. 2563 มีผู้เสียชีวิตจากการเป็นโรคมะเร็งลำไส้ใหญ่มากกว่า 930,000 ราย คิดเป็นร้อยละ 9.4 ของผู้เสียชีวิตด้วยโรคมะเร็งทั่วโลก โรคมะเร็งชนิดนี้จัดเป็นมะเร็งชนิดที่พบบ่อยเป็นอันดับสองในผู้หญิง และเป็นอันดับสามในผู้ชาย อุบัติการณ์ของโรคมะเร็งลำไส้ได้เพิ่มขึ้นอย่างต่อเนื่อง โดยพบผู้ป่วยมะเร็งลำไส้ใหญ่รายใหม่มากกว่า 1.9 ล้านราย และคิดเป็นประมาณ 10 % ของผู้ป่วยโรคมะเร็งทั้งหมด (Heisser et al., 2023; World Health Organization, 2023) องค์การอนามัยโลก (World Health Organization: WHO) ได้เผยแพร่ข้อมูลเกี่ยวกับโรคมะเร็งลำไส้ใหญ่ว่า มักได้รับการวินิจฉัยโรคในระยะลุกลาม ซึ่งมีข้อจำกัดในการรักษามากกว่าระยะเริ่มต้น (Purnama et al., 2023) การตรวจคัดกรองและวินิจฉัยโรคมะเร็งลำไส้ใหญ่อย่างทันท่วงทีมีความสำคัญต่อการรักษาและเพิ่มอัตราการรอดชีวิตของผู้ป่วย

ข้อมูลจากองค์การอนามัยโลก (WHO) พ.ศ 2563 ประเทศไทยมีผู้ป่วยโรคมะเร็งรายใหม่ 190,363 คนต่อปี หรือเพิ่มขึ้นมากกว่า 500 คนต่อวัน และพบผู้เสียชีวิตจากโรคมะเร็ง 124,866 คนต่อปี หรือมากกว่า 300 คนต่อวัน (International Agency for Research on Cancer, n.d.) นอกจากการเสียชีวิตแล้ว ยังทำให้เกิดสูญเสียทางเศรษฐกิจและเพิ่มภาระงานด้านระบบสาธารณสุขไทยเป็นมูลค่าสูง โดยในปี พ.ศ. 2562 พบว่า มูลค่าการสูญเสียทางเศรษฐกิจประมาณ 184,600 ล้านบาทต่อปี (Thitima, 2023) นอกจากนี้ ประเทศไทยพบโรคมะเร็งลำไส้ใหญ่เป็นอันดับสามในเพศชาย รองจากโรคมะเร็งตับและมะเร็ง

ปอด และพบมากเป็นอันดับสองในเพศหญิงรองจากโรคมะเร็งเต้านม โดยพบผู้ป่วยรายใหม่เป็นเพศชาย 10,660 รายต่อปี คิดเป็น 11.4% และเพศหญิง 10,443 รายต่อปี คิดเป็น 10.7% โดยไทยมีอัตราการเสียชีวิตด้วยโรคมะเร็งลำไส้ปรับมาตรฐานอายุ (Age-Standardized Mortality Rate: ASMR) อยู่ที่ 8.4 (International Agency for Research on Cancer, n.d.) เปรียบเทียบกับสัดส่วนทั่วโลกมีอัตราอยู่ที่ 9.0 (World Cancer Research Fund International, n.d.)

โรคมะเร็งลำไส้ใหญ่เกิดจากการเจริญเติบโตของเซลล์ลำไส้ใหญ่ที่ผิดปกติจนกลายเป็นมะเร็ง เซลล์มะเร็งเหล่านี้จะแบ่งตัวและแพร่กระจายไปยังส่วนต่าง ๆ ของร่างกายได้ การตรวจคัดกรองโรคมะเร็งลำไส้ใหญ่เป็นวิธีที่สำคัญในการลดอัตราการเสียชีวิตจากโรคนี้ วิธีการตรวจคัดกรองที่เป็นที่ยอมรับสามารถทำได้ 4 วิธี ได้แก่ 1) การตรวจเลือดเพื่อหาสารบ่งชี้การเป็นโรคมะเร็งลำไส้ใหญ่ (Fecal Immunochemical Test: FIT) 2) การตรวจเอกซเรย์คอมพิวเตอร์ลำไส้ใหญ่ (Computed tomography colonography: CT colonography) 3) การส่องกล้องลำไส้ใหญ่ (Colonoscopy) และ 4) การตรวจด้วยภาพทางจุลพยาธิวิทยา (Histopathologic screening) โดยในปัจจุบัน การตรวจคัดกรองด้วยภาพทางจุลพยาธิวิทยาเป็นวิธีการตรวจคัดกรองที่มีความแม่นยำสูงที่สุดแต่ต้องใช้ผู้เชี่ยวชาญในการอ่านภาพ โดยอาศัยภาพถ่ายของเนื้อเยื่อที่ตัดออกมาจากลำไส้ใหญ่เพื่อตรวจหาเซลล์มะเร็ง แพทย์จะตรวจหาลักษณะที่บ่งชี้ว่าเซลล์เป็นมะเร็ง เช่น การเจริญเติบโตของเซลล์ที่ไม่เป็นระเบียบ การสูญเสียโครงสร้างปกติของเนื้อเยื่อ และการเปลี่ยนแปลงของยีน อย่างไรก็ตาม วิธีการนี้ต้องอาศัยทักษะและประสบการณ์ของแพทย์ผู้เชี่ยวชาญในการอ่านภาพ ซึ่งอาจนำไปสู่ความคลาดเคลื่อนและความแปรปรวนในการวินิจฉัย จากการศึกษาของ Smits et al. (2022) ศึกษาความแปรปรวนของผู้สังเกตการณ์ระหว่างนักพยาธิวิทยาแต่ละคน และคณะพยาธิวิทยาในการประเมินทางจุลพยาธิวิทยาของเนื้องอกในลำไส้ใหญ่ชั้นสูงในการตรวจคัดกรองมะเร็งลำไส้ของชาวเดนมาร์ก พบว่า ผู้เชี่ยวชาญให้ความแตกต่างอย่างน้อยหนึ่งรายการใน 44 กรณีจาก 83 กรณี คิดเป็น 53.0% ซึ่งอาจส่งผลต่อการเลือกการรักษา เกิดต้นทุนในการตรวจคัดกรองสูงเกินความจำเป็น นอกจากนี้ การตรวจคัดกรองด้วยแพทย์ผู้เชี่ยวชาญอาจเกิดความล่าช้า เนื่องจากข้อจำกัดของจำนวนผู้เชี่ยวชาญต่อจำนวนผู้ป่วยที่มีแนวโน้มเพิ่มสูงขึ้น ดังนั้น ในปัจจุบันเทคโนโลยีปัญญาประดิษฐ์ (Artificial Intelligence: AI) โดยเฉพาะเทคนิคการเรียนรู้ของเครื่อง (Machine Learning: ML) ได้ถูกนำมาประยุกต์ใช้กับการคัดกรองโรคจากภาพอย่างแพร่หลาย การเรียนรู้ของเครื่องสามารถเรียนรู้ลักษณะเฉพาะของโรคจากภาพทางกายภาพและสามารถใช้ในการแยกแยะโรคออกจากกันได้อย่างมีประสิทธิภาพ

งานวิจัยที่ผ่านมาที่มีการใช้เทคนิคการเรียนรู้ของเครื่องในการคัดกรองโรคจากภาพ ได้แก่ งานวิจัยของ Gupta et al. (2019) ได้นำเสนอเทคนิคการเรียนรู้ของเครื่องเพื่อทำนายระยะของเนื้องอกในโรคมะเร็งลำไส้ใหญ่และการรอดชีวิตโดยปราศจากโรค (Disease-Free Survival: DFS) เป็นเวลา 5 ปี วิเคราะห์ข้อมูลจากผู้ป่วยมะเร็งลำไส้ใหญ่จำนวน 4,021 ราย พบว่า เทคนิคต้นไม้ป่าสุ่ม (Random Forest) มีประสิทธิภาพสูงสุดในการทำนายทั้งระยะของเนื้องอกและการรอดชีวิตโดยปราศจากโรคได้อย่างแม่นยำถึง 84% รวมถึงงานวิจัยของ Habib & Rahman (2021) ได้นำเสนอเทคนิคการเรียนรู้ของเครื่อง เพื่อตรวจคัดกรองโรค COVID-19 ด้วย 2 กระบวนการ ได้แก่ 1) การคัดกรองยีนด้วยเทคนิค eXtreme Gradient Boosting, Naïve Bayes, Regularized Random Forest, Random Forest Rule-Based Model, Random Ferns, C5.0 และ Multi-Layer Perceptron พบว่า มีความแม่นยำกว่า 93% และ 2) การ

จำแนกภาพเอกซเรย์ทรวงอกด้วยเทคนิค Deep Convolutional Neural Networks พบว่า มีความแม่นยำมากกว่า 99%

ดังนั้น งานวิจัยนี้มีวัตถุประสงค์เพื่อศึกษาการสร้างตัวแบบและเปรียบเทียบประสิทธิภาพของตัวแบบการเรียนรู้ของเครื่องสำหรับการจำแนกประเภทโรคมะเร็งลำไส้ใหญ่โดยใช้ภาพทางจุลพยาธิวิทยา การศึกษานี้จะใช้ชุดข้อมูลภาพทางจุลพยาธิวิทยาของเซลล์มะเร็งลำไส้ใหญ่และเซลล์ปกติจำนวน 10,000 ภาพ เพื่อพัฒนาตัวแบบการเรียนรู้ของเครื่อง โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพ ซึ่งทำการเปรียบเทียบทั้งหมด 4 เทคนิค ได้แก่ 1) เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machine: SVM) 2) เทคนิคเพื่อนบ้านใกล้ที่สุด (k-Nearest Neighbors: k-NN) 3) เทคนิคโครงข่ายประสาทเทียม (Neural Network: NN) และ 4) เทคนิคต้นไม้ตัดสินใจ (Decision Tree: DT) ซึ่งผลการศึกษานี้จะช่วยให้ทราบประสิทธิภาพของตัวแบบการเรียนรู้ของเครื่องในการจำแนกประเภทโรคมะเร็งลำไส้ใหญ่จากภาพทางจุลพยาธิวิทยา ซึ่งอาจนำไปสู่การประยุกต์ใช้เทคนิคการเรียนรู้ของเครื่องในการปรับปรุงประสิทธิภาพของการตรวจคัดกรองผู้ป่วยโรคมะเร็งลำไส้ใหญ่ที่สามารถช่วยแพทย์ในการวินิจฉัยโรคมะเร็งลำไส้ใหญ่ได้อย่างแม่นยำและรวดเร็ว และช่วยให้ผู้ป่วยได้รับการรักษาที่เหมาะสมและทันเวลาที่

2. วิธีดำเนินการวิจัย

การสร้างตัวแบบและเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการจำแนกประเภทโรคมะเร็งลำไส้ โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพ ผู้วิจัยได้ดำเนินการวิจัยตามกระบวนการมาตรฐานของการทำเหมืองข้อมูล (Cross-Industry Standard Process for Data Mining: CRISP-DM) (Wirth & Hipp, 2000) ซึ่งประกอบด้วย 6 ขั้นตอน ดังนี้

2.1 การทำความเข้าใจเกี่ยวกับปัญหา (Business understanding)

ผู้วิจัยได้ศึกษาข้อมูลเกี่ยวกับการเกิดโรคมะเร็งลำไส้ พบปัญหาที่เกิดขึ้นจริงเกี่ยวกับโรคมะเร็งลำไส้ ซึ่งมักไม่ทราบปัจจัยเสี่ยงของการเกิดโรคมะเร็งลำไส้ และประชาชนจำนวนมากยังขาดการเข้ารับการตรวจคัดกรอง ทำให้มักพบเมื่อป่วยเข้าสู่ระยะสุดท้าย ซึ่งแต่ละบุคคลมีพฤติกรรมการใช้ชีวิตแตกต่างกัน จากสถิติล่าสุดพบว่าการเสียชีวิตจากโรคมะเร็งลำไส้ยังมีแนวโน้มเพิ่มสูงขึ้นเรื่อย ๆ อย่างต่อเนื่อง ในปัจจุบันเทคโนโลยีทางการแพทย์และความเชี่ยวชาญของแพทย์กับการตรวจรักษาโรคมะเร็งลำไส้ใหญ่และทวารหนัก มีการพัฒนาไปอย่างมาก จึงทำให้การรักษามีประสิทธิภาพมากขึ้น โดยการตรวจทางพยาธิวิทยาเพื่อประกอบการวางแผนการรักษา ดังนั้น ผู้วิจัยจึงได้สนใจนำเทคนิคการทำเหมืองข้อมูลมาใช้ในการสร้างตัวแบบและเปรียบเทียบประสิทธิภาพของตัวแบบ สำหรับการจำแนกประเภทโรคมะเร็งลำไส้โดยใช้ภาพทางจุลพยาธิวิทยา เพื่อวางแผนการรักษาและลดปริมาณผู้ป่วยโรคมะเร็งลำไส้ในอนาคต

2.2 การทำความเข้าใจเกี่ยวกับข้อมูล (Data understanding)

ผู้วิจัยได้รวบรวมข้อมูลที่เกี่ยวข้องเพื่อใช้ในการสร้างตัวแบบและเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการจำแนกประเภทโรคมะเร็งลำไส้ ด้วยเทคนิคการทำเหมืองข้อมูลภาพ ซึ่งข้อมูลที่ใช้เป็นข้อมูลจริงที่ได้เก็บรวบรวมข้อมูลจากชุดข้อมูลภาพทางจุลพยาธิวิทยา ซึ่งเป็นภาพถ่ายจาก Pixabay ถูก

บันทึกไว้ในปี ค.ศ. 2019 ชุดข้อมูลนี้สร้างโดย Borkowski et al. (2019) ซึ่งได้เผยแพร่อย่างเปิดเผยใน
เว็บไซต์ www.kaggle.com และสามารถเข้าถึงได้อย่างเปิดเผย ชื่อชุดข้อมูล “colon_image_sets” โดย
แบ่งเป็น 2 โฟลเดอร์ ได้แก่ โฟลเดอร์ colon_aca สำหรับการจัดเก็บข้อมูลภาพทางจุลพยาธิวิทยาของ
ผู้ป่วยโรคมะเร็งลำไส้ จำนวน 5,000 ภาพ และโฟลเดอร์ colon_n สำหรับการจัดเก็บข้อมูลภาพทางจุล
พยาธิวิทยาของคนปกติ จำนวน 5,000 ภาพ รวมทั้งหมดจำนวน 10,000 ภาพ รูปภาพทั้งหมดมีขนาด 768
x 768 พิกเซล และอยู่ในรูปแบบไฟล์ .jpeg ตัวอย่างภาพทางจุลพยาธิวิทยาที่นำมาใช้ในการวิเคราะห์ดังรูป
ที่ 1



รูปที่ 1. ตัวอย่างภาพทางจุลพยาธิวิทยาที่นำมาใช้ในการวิเคราะห์

จาก Lung and Colon Cancer Histopathological Image Dataset (LC25000) (Borkowski et al., 2019)

2.3 การเตรียมข้อมูล (Data preparation)

ผู้วิจัยได้ดำเนินการเตรียมข้อมูลก่อนการวิเคราะห์ เพื่อให้เกิดความน่าเชื่อถือและข้อมูลมีคุณภาพ
เพียงพอที่จะนำมาใช้ในการวิเคราะห์ ซึ่งในงานวิจัยนี้ได้แบ่งการเตรียมข้อมูลออกเป็น 2 ขั้นตอน ดังนี้

1) การนำเข้าข้อมูลภาพมาวิเคราะห์ด้วยโปรแกรม RapidMiner Studio version 10 โดยผู้วิจัยได้
ทำการแบ่งเก็บภาพแต่ละกลุ่มไว้ในแต่ละโฟลเดอร์ 2 โฟลเดอร์ย่อย ได้แก่ โฟลเดอร์ colon_aca สำหรับ
การจัดเก็บข้อมูลภาพทางจุลพยาธิวิทยาของผู้ป่วยโรคมะเร็งลำไส้ จำนวน 5,000 ภาพ และ colon_n
สำหรับการจัดเก็บข้อมูลภาพทางจุลพยาธิวิทยาของคนปกติ จำนวน 5,000 ภาพ

2) การแปลงข้อมูลให้อยู่ในรูปแบบข้อมูลที่มีโครงสร้าง (Structured data) สำหรับงานวิจัยนี้ผู้วิจัย
ได้ใช้การแยกคุณลักษณะโดยธรรมชาติ สำหรับการทำให้เหมือนรูปภาพในระดับสากล (Global level) เพื่อ
แยกคุณสมบัติส่วนกลางออกจากภาพและบ่งบอกถึงความแตกต่างของภาพ ซึ่งเหมาะสำหรับการจำแนก
ประเภทของภาพและการกำหนดคุณสมบัติของภาพโดยรวม การคำนวณค่าต่าง ๆ จากรูปภาพแบบ
ระดับสีเทา (Grayscale) ซึ่งในงานวิจัยนี้ทำการแยกคุณสมบัติของภาพทั้งหมดโดยใช้ค่าคุณสมบัติทั้ง 7 ค่า
ได้แก่ ค่าเฉลี่ย (Mean) ค่ามัธยฐาน (Median) ค่าขั้นต่ำ (Minimum) ค่ามาตรฐานสำหรับแกน X
(Normalized X) ค่ามาตรฐานสำหรับแกน Y (Normalized Y) ค่าส่วนเบี่ยงเบนมาตรฐาน (Standard
Deviation: SD) ค่าสีเทาขั้นต่ำ (Minimum gray value) และค่าสีเทาสูงสุด (Maximum gray value) โดย
โปรแกรม RapidMiner Studio มีการติดตั้งส่วนขยายการทำเหมืองรูปภาพ (Image mining extension)
ไว้สำหรับการแยกคุณลักษณะ ตัวอย่างข้อมูลที่ถูกแปลงให้อยู่ในรูปแบบข้อมูลโครงสร้าง แสดงดังตารางที่ 1

ตารางที่ 1. ตัวอย่างข้อมูลภาพทางจุลพยาธิวิทยาที่ถูกแปลงให้อยู่ในรูปแบบของข้อมูลที่มีโครงสร้าง

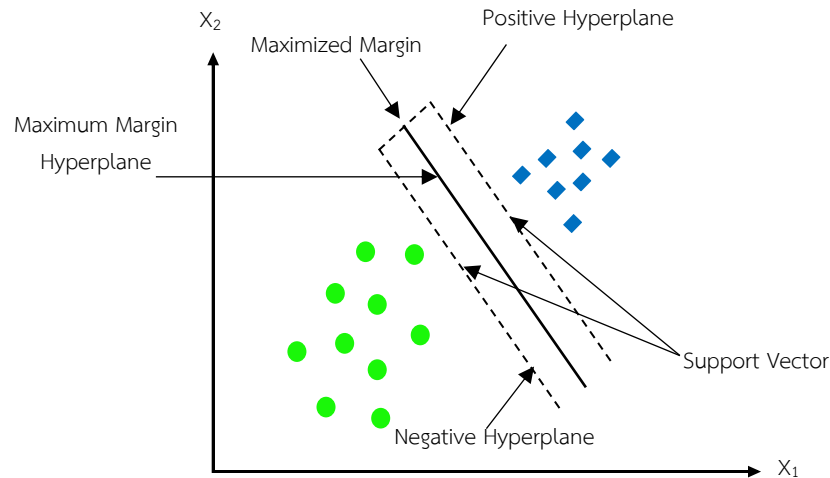
Row No.	Label	Global statistics						
		Mean	Median	Minimum Gray Value	Maximum Gray Value	Normalized X Center of Mass	Normalized Y Center of Mass	SD
1	colon_aca	216.0387	223	3	255	0.5001	0.4992	29.5406
2	colon_aca	213.7845	219	24	255	0.5051	0.5006	31.2155
3	colon_aca	218.7142	221	29	255	0.5007	0.4987	25.9928
4	colon_aca	216.4806	224	38	255	0.5039	0.5081	31.1943
5	colon_aca	204.9244	209	26	255	0.4924	0.4965	34.8963
6	colon_aca	203.9964	206	30	255	0.4982	0.5067	34.0326
7	colon_aca	206.1248	211	40	255	0.4984	0.4956	32.6657
8	colon_aca	211.2105	215	46	255	0.5017	0.5008	23.3663
9	colon_aca	213.0633	219	35	255	0.4971	0.5011	26.3662
10	colon_aca	214.0872	219	34	255	0.4930	0.4986	29.6324
:	:	:	:	:	:	:	:	:

2.4 การสร้างตัวแบบ (Modeling)

ในขั้นตอนนี้อยู่วิจัยได้นำข้อมูลที่ผ่านมากระบวนการเตรียมข้อมูลมาวิเคราะห์ด้วยเทคนิคการเรียนรู้ของเครื่อง สำหรับการสร้างตัวแบบการจำแนกประเภทโรคมะเร็งลำไส้ ด้วยเทคนิคการประมวลผลภาพ (Image processing) ในการสร้างตัวแบบเพื่อจำแนกประเภทเซลล์มะเร็งลำไส้ใหญ่ ผู้วิจัยได้เลือกเทคนิค 4 วิธีที่เป็นที่นิยมและมีการประยุกต์ใช้อย่างแพร่หลายในงานจำแนกภาพทางการแพทย์ ได้แก่ เทคนิค SVM เทคนิค k-NN เทคนิค NN และเทคนิค DT โดยพิจารณาจากคุณสมบัติของชุดข้อมูลที่มีขนาดค่อนข้างเล็กและโครงสร้างข้อมูลที่ไม่ซับซ้อนมาก ตลอดจนความสามารถของแต่ละเทคนิคในการตีความผลลัพธ์ได้ดีในบริบททางการแพทย์ โดยใช้โปรแกรม RapidMiner Studio version 10 สำหรับการสร้างตัวแบบการจำแนกประเภทข้อมูลภาพในครั้งนี้ (Sukprasert, 2023)

ในงานวิจัยนี้ได้เลือกใช้ชุดข้อมูลจากเว็บไซต์ www.kaggle.com ซึ่งเป็นแหล่งข้อมูลที่มีความนิยมอย่างแพร่หลายสำหรับการวิเคราะห์ภาพทางการแพทย์และการจำแนกโรค เนื่องจากมีความหลากหลายของข้อมูลและความพร้อมใช้งานที่เปิดเผยสาธารณะ (Almukhtar et al., 2024; Vommi & Battula, 2023; Gnanaselvi & Kalavathy, 2021; Krishnan et al., 2022; Zerouaoui & Idri, 2021) การใช้ชุดข้อมูลจากเว็บไซต์ www.kaggle.com จึงเหมาะสมกับวัตถุประสงค์ของการสร้างและเปรียบเทียบโมเดลในงานวิจัยนี้

1) เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (SVM) จัดอยู่ในกลุ่มของการเรียนรู้แบบมีผู้สอน (Supervised learning) และเป็นอัลกอริทึมที่มีประสิทธิภาพสูงในการจำแนกประเภทข้อมูล โดยเฉพาะกับชุดข้อมูลขนาดเล็กและมีความซับซ้อน เทคนิค SVM มีการนำค่าของกลุ่มข้อมูลวางในพีเจอร์สเปซ (Feature space) จากนั้นจึงหาเส้นที่ใช้ในการแบ่งข้อมูลออกจากกัน จากนั้นสร้างเส้นแบ่ง (Hyperplane) เป็นเส้นตรงขึ้นมา และหาเส้นตรงที่แบ่งข้อมูลสองกลุ่มออกจากกันที่ดีที่สุด ดังรูปที่ 2



รูปที่ 2. เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (SVM)

เทคนิค SVM เป็นเทคนิคที่ได้รับความนิยมอย่างสูงในงานจำแนกข้อมูลทางการแพทย์ โดยเฉพาะข้อมูลที่มีลักษณะเชิงซ้อนและมีชุดข้อมูลจำนวนน้อย Almkhatar et al. (2024) แสดงให้เห็นว่า SVM สามารถจำแนกภาพ MRI เพื่อวินิจฉัยเนื้องอกในสมองได้อย่างแม่นยำถึง 89% จากการใช้ชุดข้อมูลของเว็บไซต์ www.kaggle.com โดยอาศัยหลักการสร้างเส้นแบ่งข้อมูล (Hyperplane) ที่เหมาะสมที่สุดในเชิงคณิตศาสตร์ ซึ่งทำให้เทคนิคนี้เหมาะสำหรับข้อมูลที่มีการกระจายตัวที่ไม่แน่นอน เช่นเดียวกับชุดข้อมูลภาพจุลพยาธิวิทยาที่ใช้ในงานวิจัยนี้

2) เทคนิคเพื่อนบ้านใกล้ที่สุด (k-NN) เป็นวิธีการจำแนกประเภทข้อมูลวิธีหนึ่งที่ได้รับความนิยมอย่างมาก ซึ่งสามารถนำไปประยุกต์ใช้กับงานได้หลากหลาย โดยลักษณะการทำงานของเทคนิค k-NN นี้จะไม่มีการสร้างโมเดลในการจำแนกประเภทข้อมูลไว้ล่วงหน้า แต่ใช้วิธีการนำข้อมูลใหม่ (Unseen data) ที่ต้องการจำแนกมาเทียบกับข้อมูลเดิมที่คล้ายคลึงกัน และจำแนกประเภทข้อมูลใหม่โดยอาศัยการตรวจสอบค่า k แบบไขว้ (Cross validation) โดยคำนวณได้จากสมการที่ (1)

$$D_{\text{Euclidian}}(x_i, y_i) = \sqrt{\sum_{k=1}^k (x_i - y_i)^2} \quad (1)$$

โดยที่ $D_{\text{Euclidian}}(x_i, y_i)$ คือ ระยะห่างระหว่างตัวอย่าง x_i กับตัวอย่าง y_i
 k คือ คุณลักษณะทั้งหมดของตัวอย่าง (i) ที่มีค่าระยะห่างน้อยที่สุด k ตัว
 เพื่อนำมาพิจารณาคำตอบ

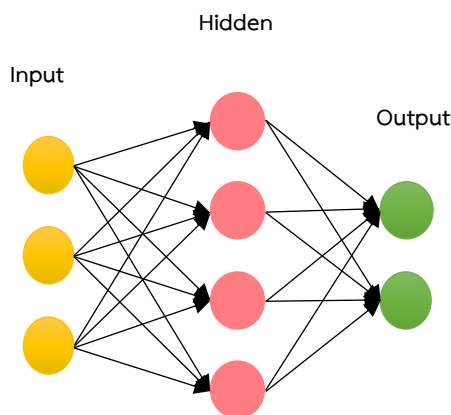
เทคนิค k-NN เป็นเทคนิคที่ง่ายและมีความสามารถในการจำแนกข้อมูลภาพทางการแพทย์ที่มีโครงสร้างชัดเจนและข้อมูลไม่ซับซ้อนมาก Vommi & Battula (2023) ได้พัฒนา Fuzzy k-NN เพื่อจำแนกข้อมูล COVID-19 จากชุดข้อมูลเว็บไซต์ www.kaggle.com และพบว่ามีประสิทธิภาพสูงกว่า k-NN แบบดั้งเดิมอย่างมีนัยสำคัญ ความสำเร็จของ k-NN มักเกิดจากความสามารถในการจำแนกข้อมูลที่มีขนาดชุดข้อมูลเล็ก (Small sample size) และข้อมูลที่มีขอบเขตจำแนกชัดเจน (Clear decision boundary) โดยไม่ต้องมีการตั้งสมมติฐานล่วงหน้ากับข้อมูล (Non-parametric assumption) ซึ่งตรงกับลักษณะของข้อมูลภาพจุลพยาธิวิทยาในงานวิจัยนี้ที่มีขนาดชุดข้อมูลไม่ใหญ่มาก และมีลักษณะแบ่งกลุ่มระหว่างเซลล์มะเร็งกับเซลล์ปกติที่ค่อนข้างชัดเจน

3) เทคนิคโครงข่ายประสาทเทียม (NN) เป็นวิธีการของโครงข่ายประสาทที่ให้เครื่องได้เรียนรู้จากตัวอย่างต้นแบบ และฝึกให้ระบบได้รู้จักคิดแก้ปัญหาที่กว้างขึ้นได้ โดยในโครงสร้างของโครงข่ายประสาทประกอบด้วย โหนดสำหรับข้อมูลเข้า-ข้อมูลออก (Input-Output data) และการประมวลผลกระจายอยู่ในโครงสร้างเป็นชั้น ๆ ดังรูปที่ 3 และใช้การคำนวณจากสมการที่ (2)

$$w_{ij}x_{ij} = w_{0j}x_{0j} + \dots + w_{ij}x_{ij} \quad (2)$$

โดยที่ x_{ij} คือ ข้อมูลเข้าที่ i ไปยังโหนดที่ j

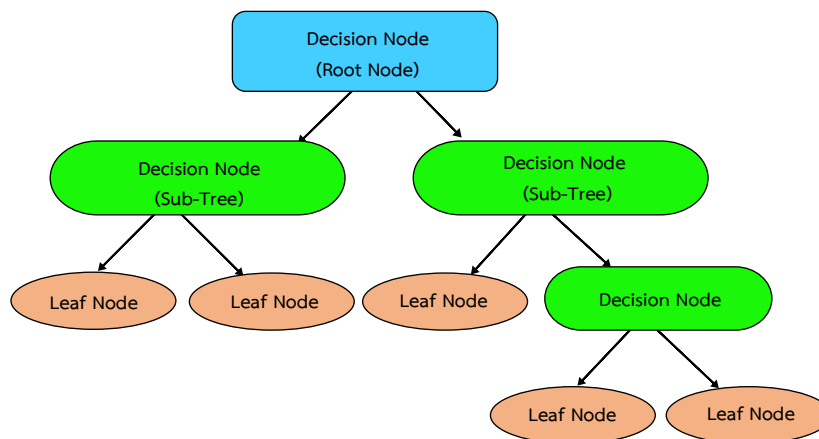
w_{ij} คือ น้ำหนักถ่วงที่สัมพันธ์กับข้อมูลเข้าที่ i ไปยังโหนดที่ j และมีข้อมูลเข้าจำนวน $i + 1$ ไปยังโหนด j



รูปที่ 3. เทคนิคโครงข่ายประสาทเทียม

เทคนิคโครงข่ายประสาทเทียมโดยเฉพาะ Convolutional Neural Networks (CNN) ได้รับการพิสูจน์ว่าให้ประสิทธิภาพสูงในการจำแนกภาพทางการแพทย์ที่มีความซับซ้อนและมีข้อมูลจำนวนมาก Gnanaselvi & Kalavathy (2021) รายงานว่า CNN สามารถจำแนกความผิดปกติในภาพจอประสาทตาได้อย่างแม่นยำสูง โดยสามารถเรียนรู้คุณลักษณะเชิงลึก (Deep features) จากข้อมูลภาพโดยอัตโนมัติ ดังนั้นแม้ว่างานวิจัยนี้จะมีขนาดข้อมูลไม่ใหญ่มาก แต่การทดลองใช้ NN ช่วยเปิดโอกาสในการวิเคราะห์ลักษณะเชิงลึกของเซลล์ในภาพได้

4) เทคนิคต้นไม้ตัดสินใจ (DT) เป็นอัลกอริธึมในการจำแนกประเภทข้อมูลที่ใช้กับข้อมูลแบบต่อเนื่องและไม่ต่อเนื่อง โดยใช้หลักการต้นไม้ในการค้นหาคุณลักษณะที่ดีที่สุดจากบนลงล่าง เพื่อสร้างเป็นโหนดราก (Root node) ใช้ค่า Gain Ratio ที่สูงที่สุดเป็นโหนดรากและโหนดถัดไป โครงสร้างเทคนิค DT แสดงดังรูปที่ 4 ค่า Entropy Information Gain และ Split Information (Xu et al., 2013) ใช้ในการหาค่า Information Gain (IG) โดยคำนวณได้จากสมการที่ (3)



รูปที่ 4. โครงสร้างเทคนิคต้นไม้ตัดสินใจ

$$IG(\text{parent}, \text{child}) = \text{entropy}(\text{parent}) - [p(c_1) \times \text{entropy}(c_1) + p(c_2)\text{entropy}(c_2) + \dots] \quad (3)$$

โดยที่ $\text{entropy}(c_1) = -p(c_1) \log_2 p(c_1)$ และ $p(c_1)$ คือ ค่าความน่าจะเป็นของ c_1

เทคนิค DT เป็นเทคนิคที่มีความสามารถในการตีความโมเดลได้ง่ายและเหมาะสำหรับการวิเคราะห์ข้อมูลขนาดเล็ก Krishnan et al. (2022) แสดงให้เห็นว่าการใช้ DT และ Random Forest สามารถจำแนกข้อมูลทางการแพทย์ได้อย่างมีประสิทธิภาพในบริบทที่ต้องการความโปร่งใสในการตีความ (Interpretability) ซึ่งเป็นประเด็นสำคัญในการวิเคราะห์ข้อมูลเชิงคลินิกเช่นเดียวกับในงานวิจัยนี้

การประเมินประสิทธิภาพของตัวแบบ ในงานวิจัยนี้ได้ใช้วิธีการตรวจสอบความถูกต้องข้ามแบบ 10 เท่า (10-Fold cross-validation) ควบคู่กับการประมวลผลซ้ำ 30 รอบ เพื่อเพิ่มความเสถียรและลด

อิทธิพลของการสุ่มในการแบ่งข้อมูล โดยแนวทางนี้สอดคล้องกับข้อเสนอแนะของ Górriz et al. (2024) และ Krstajic et al. (2014) ที่ระบุว่าการทำ Cross-validation ซ้ำหลายครั้งช่วยลดความแปรปรวน (Variance) และให้การประเมินผลลัพธ์ของตัวแบบที่น่าเชื่อถือยิ่งขึ้น นอกจากนี้ ยังช่วยลดโอกาสการเกิด Overfitting ได้อย่างมีนัยสำคัญ (Shi & Adnan, 2014; Wilimitis & Walsh, 2023)

การเลือกจำนวน 30 ครั้ง เป็นแนวทางที่เหมาะสมกับการประเมินตัวแบบบนชุดข้อมูลที่มีลักษณะกระจายไม่สมมาตรและมีการสุ่มในการแบ่งข้อมูลหลายรอบ เพื่อลดความเียงเบนของค่าเฉลี่ย (Bias) และเพิ่มความแม่นยำของการประเมิน (Raschka, 2018) นอกจากนี้ เพื่อประเมินปัญหา Overfitting และ Underfitting จึงได้ทำการวิเคราะห์ค่าความแปรปรวนด้วยการทดสอบทางสถิติ เช่น Welch's Test และ Levene's Test เพื่อตรวจสอบหาความแตกต่างของประสิทธิภาพระหว่างตัวแบบแต่ละตัวมีนัยสำคัญทางสถิติหรือไม่ แต่ไม่ได้ใช้เพื่อวินิจฉัยการเกิด Overfitting หรือ Underfitting โดยตรง (Raschka, 2018; Markatou et al., 2005)

ผลการวิเคราะห์แสดงให้เห็นว่าแม้ Welch's Test และ Levene's Test จะให้ค่า p-value ต่ำกว่า 0.01 (เช่น $p = 0.000$) หมายความว่า ความแตกต่างอย่างมีนัยสำคัญทางสถิติระหว่างโมเดล แต่ไม่ได้บ่งชี้ว่าตัวแบบใดตัวแบบหนึ่งมีปัญหา Overfitting อย่างชัดเจน การใช้เทคนิค Repeated cross-validation และการเปรียบเทียบผลลัพธ์ค่าเฉลี่ยหลายรอบ จึงเป็นแนวทางที่เพียงพอในการควบคุมและลดความเสี่ยงของการเกิด Overfitting ภายในงานวิจัยนี้

ในการประเมินประสิทธิภาพของแต่ละเทคนิค ผู้วิจัยได้ทำการปรับค่าพารามิเตอร์ของแต่ละโมเดลให้เหมาะสมก่อนการทดสอบ โดยได้ใช้วิธี Grid search และการทดลองเปรียบเทียบค่าพารามิเตอร์ที่แตกต่างกัน เพื่อหา Configuration ที่ให้ประสิทธิภาพดีที่สุด รายละเอียดการปรับค่าพารามิเตอร์แสดงในตารางที่ 2

ตารางที่ 2. รายละเอียดการปรับค่าพารามิเตอร์ของแต่ละเทคนิค

เทคนิค	พารามิเตอร์ที่ทดสอบ	ค่าที่ใช้ในการทดลอง	ค่าที่ได้ผลดีที่สุด
SVM	kernel, C, gamma	kernel: linear, polynomial, RBF	kernel = RBF
		C: 0.1, 1, 10	C = 1
		gamma: scale, auto	gamma = scale
k-NN	จำนวนเพื่อนบ้าน (k), ระยะทาง	k: 3, 5, 7, 9, 11, 13, 15	k = 5
		metric: Euclidean, Manhattan	metric = Euclidean
NN	hidden layers, nodes/layer, activation function	hidden layers: 1–3	2 layers
		nodes/layer: 8, 16, 32	16 nodes/layer
		activation: ReLU, Tanh	activation = ReLU
DT	splitting criterion, max depth	criterion: gini, entropy	criterion = gini
		max depth: 3, 5, 10, None	max depth = 5

ค่าพารามิเตอร์ที่พิจารณา เช่น 1) เทคนิค SVM ได้ทดสอบ kernel หลายแบบ ได้แก่ linear, polynomial, และ RBF พร้อมการปรับค่า C และ gamma 2) เทคนิค k-NN ได้ทดสอบค่าจำนวนเพื่อนบ้านตั้งแต่ 3-15 3) เทคนิค NN ได้ปรับจำนวนชั้นซ่อนและจำนวนโหนดในแต่ละชั้น รวมถึง activation function และ 4) เทคนิค DT ได้ทดลองการแบ่งข้อมูลโดยใช้เกณฑ์ gini และ entropy ร่วมกับการกำหนดระดับความลึก (Max depth) ที่แตกต่างกัน

2.5 การประเมินผล (Evaluation)

ผู้วิจัยได้ทำการแบ่งข้อมูลออกเป็น 2 ส่วน คือ ส่วนที่ใช้สำหรับการสร้างตัวแบบ (Training set) เพื่อทำหน้าที่เป็นการสร้างตัวแบบในการคัดกรองโรคมะเร็งลำไส้ และส่วนที่ใช้สำหรับการทดสอบ (Testing Set) เพื่อทำการทดสอบประสิทธิภาพของตัวแบบ ด้วยวิธี Cross validation โดยแบ่งข้อมูลออกเป็น 10 ส่วนเท่า ๆ กัน (10-Fold cross validation) และทำการทดสอบประสิทธิภาพของการจำแนกประเภทข้อมูลภาพ โดยสร้างเป็นเมทริกซ์ความสับสน (Confusion matrix) เพื่อการเก็บข้อมูลที่เกี่ยวข้องกับการแบ่งแยกข้อมูลจริงกับข้อมูลที่เกิดจากการทำนายด้วยระบบการแบ่งแยก (Classification system) ดังแสดงในตารางที่ 3 และค่าประสิทธิภาพของการจำแนกประเภทข้อมูลภาพที่ใช้งานวิจัยนี้ ได้แก่ ค่าความแม่นยำ (Accuracy) ค่าความไว (Sensitivity) ค่าจำเพาะ (Specificity) และค่าประสิทธิภาพโดยรวม (F-measure) (Sukprasert, 2023)

ตารางที่ 3. เมทริกซ์ความสับสน (Confusion matrix)

	Predicted Positive (Yes)	Predicted Negative (No)
Actual Positive (Yes)	True Positive (TP)	False Negative (FN)
Actual Negative (No)	False Positive (FP)	True Negative (TN)

โดยที่ True Positive (TP) แทนจำนวนข้อมูลที่ทำนายว่าเป็น โรค และสิ่งที่เกิดขึ้น คือ เป็นโรค
True Negative (TN) แทนจำนวนข้อมูลที่ทำนายว่า ไม่เป็นโรค และสิ่งที่เกิดขึ้น คือ ไม่เป็นโรค
False Positive (FP) แทนจำนวนข้อมูลที่ทำนายว่า เป็นโรค แต่สิ่งที่เกิดขึ้น คือ ไม่เป็นโรค
False Negative (FN) แทนจำนวนข้อมูลที่ทำนายว่า ไม่เป็นโรค แต่สิ่งที่เกิดขึ้น คือ เป็นโรค
ค่าความแม่นยำ คือ ค่าที่ตัวแบบสามารถพยากรณ์ข้อมูลผู้ป่วยที่เกิดโรคและไม่เกิดโรคได้อย่างถูกต้องต่อข้อมูลทั้งหมด คำนวณได้จากสมการที่ (4)

$$\text{Accuracy} = \left(\frac{TP+TN}{TP+TN+FP+ FN} \right) \quad (4)$$

ค่าความไว คือ ค่าที่ตัวแบบสามารถพยากรณ์ข้อมูลผู้ป่วยที่เกิดโรคได้อย่างถูกต้องต่อผู้ป่วยที่เกิดโรคจริง คำนวณได้จากสมการที่ (5)

$$\text{Sensitivity} = \frac{TP}{TP + FN} \quad (5)$$

ค่าจำเพาะ คือ ค่าที่ตัวแบบสามารถพยากรณ์ข้อมูลผู้ป่วยที่ไม่เกิดโรคได้อย่างถูกต้องต่อผู้ป่วยที่ไม่เกิดโรคจริง คำนวณได้จากสมการที่ (6)

$$\text{Specificity} = \frac{TN}{TN + FP} \quad (6)$$

ค่าประสิทธิภาพโดยรวม คือ ค่าที่เกิดจากการเปรียบเทียบระหว่าง ค่า Precision และ ค่า Recall ของแต่ละคลาสเป้าหมาย คำนวณได้จากสมการที่ (7)

$$\text{F-measure} = \frac{(2 * \text{Precision} * \text{Recall})}{(\text{Precision} + \text{Recall})} \quad (7)$$

2.6 การนำไปใช้งาน (Deployment)

เมื่อทำการวิเคราะห์ข้อมูลเพื่อหาตัวแบบที่มีความเหมาะสมสำหรับการจำแนกประเภทผู้ป่วยโรคมะเร็งลำไส้ ตัวแบบที่ได้สามารถนำไปพัฒนาเป็นระบบสารสนเทศสำหรับการคัดกรองผู้ป่วยโรคมะเร็งลำไส้ด้วยภาพทางจุลพยาธิ และสนับสนุนการตัดสินใจสำหรับแพทย์ในการวินิจฉัยคัดกรองผู้ป่วยโรคมะเร็งลำไส้ ทำให้การวินิจฉัยมีความแม่นยำและรวดเร็วยิ่งขึ้น นอกจากนี้ยังช่วยแพทย์และทีมสห-วิชาชีพในการกำหนดนโยบายด้านการวางแผนการรักษา และจัดเตรียมเวชภัณฑ์ให้เพียงพอต่อจำนวนผู้ป่วยที่จะเกิดขึ้น

3. ผลการวิจัยและอภิปรายผล

ผลการวิเคราะห์ข้อมูลการเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการจำแนกประเภทโรคมะเร็งลำไส้ด้วยเทคนิคการจำแนกประเภทข้อมูลภาพ ผู้วิจัยได้ทำการทดลองเพื่อประเมินประสิทธิภาพของตัวแบบ โดยพิจารณาค่าชี้วัดด้านประสิทธิภาพ ได้แก่ ค่าความแม่นยำ ค่าประสิทธิภาพโดยรวม ค่าความไว และค่าจำเพาะ ซึ่งดำเนินการประมวลผลตัวแบบซ้ำจำนวน 30 รอบ จากนั้นจึงวิเคราะห์หาค่าเฉลี่ย ($\bar{x}$) และส่วนเบี่ยงเบนมาตรฐาน (SD) ของค่าประสิทธิภาพการจำแนกประเภทข้อมูลภาพของแต่ละตัวแบบสำหรับการคัดกรองผู้ป่วยโรคมะเร็งลำไส้ ผลการวิเคราะห์แสดงดังตารางที่ 4 ต่อมาผู้วิจัยได้ทำการเปรียบเทียบค่าเฉลี่ยของค่าความแม่นยำของตัวแบบ เพื่อทดสอบสมมติฐานว่าเทคนิคการจำแนกประเภทข้อมูลภาพที่ใช้ในงานวิจัยนี้มีประสิทธิภาพในการจำแนกประเภทข้อมูลแตกต่างกันหรือไม่ โดยกำหนดระดับนัยสำคัญทางสถิติที่ 0.01 ผลการทดสอบสมมติฐานแสดงดังตารางที่ 5

ตารางที่ 4. ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของประสิทธิภาพการจำแนกประเภทข้อมูลภาพของทั้ง 4 เทคนิค

Classification Performance	Image Classification Techniques							
	Support Vector Machine		k-Nearest Neighbors		Neural Network		Decision Tree	
	$\bar{x}$	SD	$\bar{x}$	SD	$\bar{x}$	SD	$\bar{x}$	SD
Accuracy	79.63	0.0310	91.86	0.0000	83.99	0.0942	79.72	0.1773
Sensitivity	87.44	0.0132	92.24	0.0000	87.79	0.3103	96.36	0.3147
Specificity	71.83	0.0570	91.48	0.0000	80.02	0.4843	63.07	0.4108
F-measure	81.10	0.0263	91.89	0.0000	84.50	0.0490	82.62	0.1494

จากตารางที่ 4 พบว่า ผลการวิเคราะห์ค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของประสิทธิภาพการจำแนกประเภทข้อมูลภาพของทั้ง 4 เทคนิค พบว่า เทคนิค k-NN ให้ค่าเฉลี่ยความแม่นยำสูงสุด เท่ากับ 91.86% รองลงมา ได้แก่ เทคนิค Neural Network มีค่าเท่ากับ 83.99% ในขณะที่เทคนิค Support Vector Machine (SVM) ให้ค่าเฉลี่ยความแม่นยำต่ำที่สุด เท่ากับ 79.63% เมื่อพิจารณาค่าความไวพบว่า เทคนิค Decision Tree ให้ค่าเฉลี่ยสูงสุด เท่ากับ 96.36% รองลงมา คือ เทคนิค k-NN มีค่าเท่ากับ 92.24% และเทคนิคที่ให้ค่าเฉลี่ยความไวต่ำที่สุด คือ SVM ซึ่งมีค่าเท่ากับ 87.44% สำหรับค่าความจำเพาะพบว่า เทคนิค k-NN ให้ค่าเฉลี่ยสูงสุด เท่ากับ 91.48% รองลงมา คือ เทคนิค Neural Network มีค่าเท่ากับ 80.02% ขณะที่เทคนิค Decision Tree ให้ค่าเฉลี่ยความจำเพาะต่ำที่สุด เท่ากับ 63.07% ในส่วนของค่าประสิทธิภาพโดยรวม พบว่า เทคนิค k-NN ให้ค่าเฉลี่ยสูงสุด เท่ากับ 91.89% รองลงมา คือ เทคนิค Neural Network มีค่าเท่ากับ 84.50% และเทคนิคที่ให้ค่าเฉลี่ยค่าประสิทธิภาพโดยรวมต่ำที่สุด คือ SVM ซึ่งมีค่าเท่ากับ 81.10% ตามลำดับ

ตารางที่ 5. ผลการทดสอบสมมติฐานความแปรปรวนด้วย Levene's test และ Welch's ANOVA

Levene Statistic	Sig.	Welch test	Sig.
163.741	0.000*	1624329.683	0.000*

*มีนัยสำคัญทางสถิติที่ระดับ 0.01

ตารางที่ 6. ผลการเปรียบเทียบประสิทธิภาพการจำแนกประเภทข้อมูลภาพแบบรายคู่ด้วยวิธี Dunnett T3

(I) Classification Techniques	(J) Classification Techniques	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Support Vector Machine	k-Nearest Neighbors	-12.22533	0.02623	0.000*	-12.3098	-12.1409
	Neural Network	-4.27067	0.02623	0.000*	-4.3551	-4.1862
	Decision Tree	-0.08067	0.02623	0.016	7.8702	0.0038

ตารางที่ 6. (ต่อ) ผลการเปรียบเทียบประสิทธิภาพการจำแนกประเภทข้อมูลภาพแบบรายคู่ด้วยวิธี Dunnett T3

(I) Classification Techniques	(J) Classification Techniques	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
k-Nearest Neighbors	Neural Network	7.95467	0.02623	0.000*	12.0602	8.0391
	Decision Tree	12.14467	0.02623	0.000*	-8.8245	12.2291
Neural Network	Decision Tree	4.19000	0.02623	0.000*	4.1056	4.2744

*มีนัยสำคัญทางสถิติที่ระดับ 0.01

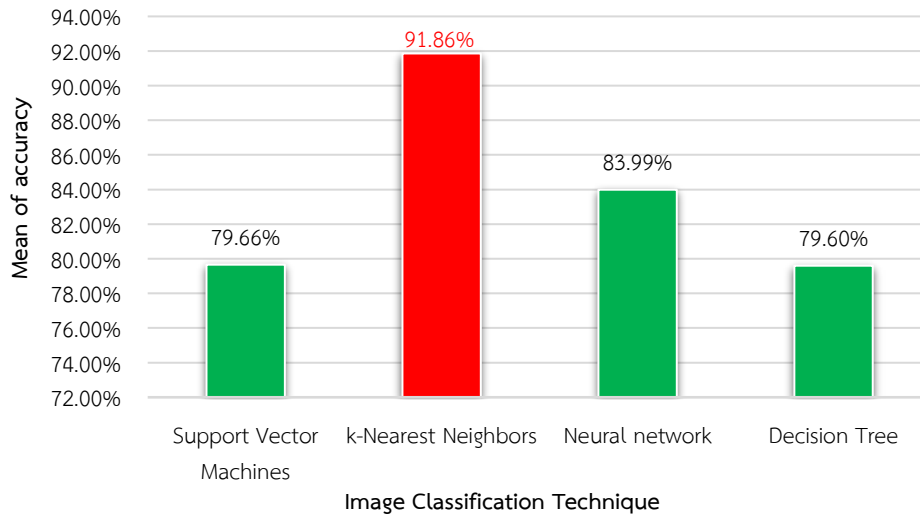
จากตารางที่ 5 แสดงผลการตรวจสอบสมมติฐานความแปรปรวนของค่าเฉลี่ยค่าความแม่นยำของเทคนิคการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค ด้วยสถิติ Levene's test พบว่า ค่า Levene statistic มีค่าเท่ากับ 163.741 และมีค่า Sig. เท่ากับ 0.000 ซึ่งน้อยกว่าระดับนัยสำคัญทางสถิติที่กำหนดไว้ที่ 0.01 แสดงว่า ความแปรปรวนของค่าเฉลี่ยค่าความแม่นยำของแต่ละเทคนิคมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ ดังนั้น ผู้วิจัยจึงเลือกใช้สถิติ Welch's ANOVA ซึ่งเหมาะสมกับกรณีที่มีสมมติฐานความแปรปรวนเท่ากันไม่เป็นจริง เพื่อทดสอบความแตกต่างของค่าเฉลี่ยค่าความแม่นยำระหว่างเทคนิคการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค ผลการทดสอบพบว่า ค่า Welch test มีค่าเท่ากับ 1,624,329.683 และมีค่า Sig. เท่ากับ 0.000 ซึ่งน้อยกว่าระดับนัยสำคัญทางสถิติที่ 0.01 แสดงให้เห็นว่า ค่าเฉลี่ยค่าความแม่นยำของเทคนิคการจำแนกประเภทข้อมูลภาพแต่ละเทคนิคมีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติ

จากผลการทดสอบดังกล่าว ผู้วิจัยจึงดำเนินการทดสอบความแตกต่างของค่าเฉลี่ยแบบรายคู่ (Post hoc test) ด้วยวิธี Dunnett T3 เพื่อเปรียบเทียบความแตกต่างระหว่างเทคนิคการจำแนกประเภทข้อมูลภาพเป็นรายคู่ ซึ่งแสดงผลดังตารางที่ 6

จากตารางที่ 6 แสดงผลการทดสอบเปรียบเทียบค่าเฉลี่ยของค่าความแม่นยำของเทคนิคการจำแนกประเภทข้อมูลภาพแบบรายคู่ ด้วยวิธี Dunnett T3 พบว่า เทคนิคส่วนใหญ่มีค่าเฉลี่ยของค่าความแม่นยำแตกต่างกันอย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 อย่างไรก็ตาม เมื่อพิจารณาเปรียบเทียบระหว่างเทคนิค Support Vector Machine (SVM) และเทคนิค Decision Tree พบว่า ค่า Sig. มีค่าเท่ากับ 0.016 ซึ่งมากกว่าระดับนัยสำคัญที่กำหนดไว้ที่ 0.01 แสดงให้เห็นว่า ค่าเฉลี่ยของค่าความแม่นยำของทั้งสองเทคนิคไม่แตกต่างกันอย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01

จากนั้น ผู้วิจัยได้นำค่าเฉลี่ยของค่าความแม่นยำของแต่ละเทคนิคการจำแนกประเภทข้อมูลภาพมาสร้างเป็นแผนภูมิแท่ง เพื่อเปรียบเทียบประสิทธิภาพของตัวแบบการจำแนกประเภทข้อมูลภาพดังแสดงในรูปที่ 5

Performance Comparison of Image Classification Models



รูปที่ 5. การเปรียบเทียบค่าเฉลี่ยของค่าความแม่นยำของตัวแบบการจำแนกประเภทข้อมูลภาพ

จากรูปที่ 5 แสดงการเปรียบเทียบค่าเฉลี่ยของค่าความแม่นยำของตัวแบบการจำแนกประเภทข้อมูลภาพ พบว่า เทคนิค k-NN ให้ค่าเฉลี่ยความแม่นยำสูงที่สุด เท่ากับ 91.86% รองลงมา คือ เทคนิค Neural Network ซึ่งมีค่าเฉลี่ยความแม่นยำเท่ากับ 83.99% ขณะที่เทคนิค Decision Tree ให้ค่าเฉลี่ยความแม่นยำต่ำที่สุด เท่ากับ 79.66%

จากผลการเปรียบเทียบดังกล่าว สามารถสรุปได้ว่า เทคนิค k-NN มีประสิทธิภาพในการจำแนกประเภทข้อมูลภาพสูงที่สุดเมื่อพิจารณาจากค่าความแม่นยำ และมีความเหมาะสมที่สุดสำหรับนำไปพัฒนาเป็นตัวแบบสำหรับการคัดกรองผู้ป่วยโรคมะเร็งลำไส้ในงานวิจัยนี้

4. สรุปผลการวิจัย

การสร้างตัวแบบและการเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการคัดกรองผู้ป่วยโรคมะเร็งลำไส้ ด้วยเทคนิคการวิเคราะห์ข้อมูลภาพ โดยใช้ภาพถ่ายทางจุลพยาธิวิทยาที่นำวิเคราะห์ตามกระบวนการมาตรฐานการทำเหมืองข้อมูล โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค ได้แก่ เทคนิค k-NN เทคนิค DT เทคนิค NN และเทคนิค SVM จากนั้นทำการเปรียบเทียบประสิทธิภาพของการจำแนกประเภทข้อมูลภาพด้วยค่าความแม่นยำ ค่าประสิทธิภาพโดยรวม ค่าความไว และค่าจำเพาะ เพื่อศึกษาผลการวิเคราะห์ของประสิทธิภาพของตัวแบบแต่ละเทคนิคว่าตัวแบบใดเหมาะสำหรับการนำไปใช้สำหรับการจำแนกประเภทโรคมะเร็งลำไส้มากที่สุด โดยใช้โปรแกรม RapidMiner Studio Version 10 สำหรับการเตรียมข้อมูล การสร้างตัวแบบและการทดสอบประสิทธิภาพของตัวแบบ ซึ่งผลการศึกษาพบว่า เทคนิค k-NN ให้ค่าความแม่นยำสูงที่สุดมีค่าเท่ากับ 91.86% ซึ่งสอดคล้องกับงานวิจัยของ Xu et al. (2013) ที่ได้

ศึกษาเกี่ยวกับการจำแนกประเภทสำหรับมะเร็งลำไส้โดยใช้ภาพทางจุลพยาธิวิทยา โดยใช้เทคนิค SVM ที่มีค่าความแม่นยำเท่ากับ 70.8% และงานวิจัยของ Verma et al. (2017) ที่ได้ศึกษาเกี่ยวกับการวิเคราะห์และการระบุนิวไต์โดยใช้เทคนิค k-NN และเทคนิค SVM ซึ่งผลการวิจัยพบว่า เทคนิค k-NN ให้ความแม่นยำสูงที่สุด โดยมีค่าเท่ากับ 89% และงานวิจัยของ Lungu et al. (2016) ที่ได้ศึกษาการวินิจฉัยความดันโลหิตสูงในปอดจากการตรวจด้วยคลื่นสนามแม่เหล็กด้วยแบบการคำนวณตามภาพและการวิเคราะห์แผนผังการตัดสินใจ ซึ่งเทคนิค DT ให้ความแม่นยำสูงที่สุดมีค่าเท่ากับ 92% และงานวิจัยของ Chinpanthana (2022) ได้ศึกษาเกี่ยวกับตัวแบบการเรียนรู้ท่าทางของมนุษย์จากการเคลื่อนไหว โดยใช้เทคนิค NN แบบสังวัตนาการและแบบหน่วยความจำระยะสั้นแบบยาว ซึ่งผลการวิจัย พบว่า เทคนิค NN ให้ความแม่นยำสูงที่สุดมีค่าเท่ากับ 81%

อย่างไรก็ตาม งานวิจัยนี้มีข้อจำกัดบางประการที่ควรพิจารณา ได้แก่ ชุดข้อมูลที่ใช้มาจากแหล่งเดียวและเป็นชุดข้อมูลที่ได้รับการจัดเตรียมไว้แล้ว ซึ่งอาจไม่สะท้อนความหลากหลายของข้อมูลในสภาพแวดล้อมจริงทางคลินิก นอกจากนี้ ตัวแบบยังไม่ได้รับการทดสอบกับข้อมูลจริงจากสถานพยาบาล ซึ่งในความเป็นจริง อาจมีความแปรปรวนของคุณภาพของภาพและบริบทการวินิจฉัยที่แตกต่างกันออกไป ดังนั้น การศึกษาต่อยอดในอนาคตควรพิจารณาใช้ชุดข้อมูลจากหลายแหล่งที่มีความหลากหลายมากขึ้น และทำการประเมินตัวแบบภายใต้สภาพแวดล้อมทางคลินิกจริง เพื่อเพิ่มความแม่นยำและความน่าเชื่อถือในการประยุกต์ใช้งานจริง

การวินิจฉัยทางการแพทย์มีความสำคัญอย่างมากสำหรับการจำแนกประเภทของโรคต่าง ๆ และผลการวิเคราะห์ของงานวิจัยนี้สามารถนำไปพัฒนาต่อยอดเป็นสารสนเทศทางการแพทย์ได้ เนื่องจากต้องมีการจำแนกข้อมูลและวิเคราะห์ก่อนที่จะนำข้อมูลส่งให้กับแพทย์ เพื่อไปทำการวินิจฉัยต่อไป อีกทั้งหากมีผู้สนใจในการศึกษาต่อยอดการคัดกรองผู้ป่วยโรคมะเร็งลำไส้ โดยใช้ชุดข้อมูลหรือตัวแปรต้นที่แตกต่างกันออกไป หรือเทคนิคการทำเหมืองข้อมูลที่นอกเหนือจากงานวิจัยนี้ รวมถึงการหาชุดข้อมูลภาพที่เหมาะสมในการวิเคราะห์ถึงปัญหา ในส่วนนี้ได้มีการเปรียบเทียบข้อมูลด้านอื่น ๆ อาจจะทำให้ตัวเลขความคาดเคลื่อนเพิ่มขึ้นหรือลดลงตามชนิดของข้อมูลได้

เอกสารอ้างอิง (References)

- Almukhtar, F. H., Kareem, S. W., & Khoshaba, F. S. (2024). Design and development of an effective classifier for medical images based on machine learning and image segmentation. *Egyptian Informatics Journal*, 25, Article 100454. <https://doi.org/10.1016/j.eij.2024.100454>
- Borkowski, A. A., Bui, M. M., Thomas, L. B., Wilson, C. P., DeLand, L. A., & Mastorides, S. M. (2019). *Lung and colon cancer histopathological image dataset (LC25000)*. arXiv. <https://arxiv.org/pdf/1912.12142>
- Chinpanthana, N. (2022). Learning model of human body movement using convolutional neural network and long short-term memory. *Journal of Information Science and Technology*, 12(1), 27-36. <https://doi.org/10.14456/jist.2022.3>

- Gnanaselvi, J. A., & Kalavathy, G. M. (2021). Detecting disorders in retinal images using machine learning techniques. *Journal of Ambient Intelligence and Humanized Computing*, 12(5), 4593-4602. <https://doi.org/10.1007/s12652-020-01841-2>
(Retraction published 2022, *Journal of Ambient Intelligence and Humanized Computing*, 14, 551)
- Górriz, J. M., Segovia, F., Ramírez, J., Ortiz, A., & Suckling, J. (2024). *Is K-fold cross validation the best model selection method for machine learning?*. arXiv.
<https://doi.org/10.48550/arXiv.2401.16407>
- Gupta, P., Chiang, S.-F., Sahoo, P. K., Mohapatra, S. K., You, J.-F., Onthoni, D. D., Hung, H.-Y., Chiang, J.-M., Huang, Y., & Tsai, W.-S. (2019). Prediction of colon cancer stages and survival period with machine learning approach. *Cancers*, 11(12), Article 2007.
<https://doi.org/10.3390/cancers11122007>
- Habib, N., & Rahman, M. M. (2021). Diagnosis of corona diseases from associated genes and X-ray images using machine learning algorithms and deep CNN. *Informatics in Medicine Unlocked*, 24, Article 100621. <https://doi.org/10.1016/j.imu.2021.100621>
- Heisser, T., Hoffmeister, M., & Brenner, H. (2023). Colorectal cancer: A health and economic problem. *Best Practice & Research Clinical Gastroenterology*, 67, Article 101839. <https://doi.org/10.1016/j.bpg.2023.101839>
- International Agency for Research on Cancer. (n.d.). *Thailand fact sheets: Cancer today*. International Agency for Research on Cancer.
<https://gco.iarc.fr/today/data/factsheets/populations/764-thailand-fact-sheets.pdf>
- Krishnan, S. N., Barua, S., Frankel, T. L., & Rao, A. (2022). Towards the characterization of the tumor microenvironment through dictionary learning-based interpretable classification of multiplexed immunofluorescence images. *Physics in Medicine & Biology*, 68(1), Article 014002. <https://doi.org/10.1088/1361-6560/aca86a>
- Krstajic, D., Buturovic, L. J., Leahy, D. E., & Thomas, S. (2014). Cross-validation pitfalls when selecting and assessing regression and classification models. *Journal of Cheminformatics*, 6(1), Article 10. <https://doi.org/10.1186/1758-2946-6-10>
- Lungu, A., Swift, A. J., Capener, D., Kiely, D., Hose, R., & Wild, J. M. (2016). Diagnosis of pulmonary hypertension from magnetic resonance imaging-based computational models and decision tree analysis. *Pulmonary Circulation*, 6(2), 181-190.
<https://doi.org/10.1086/686020>
- Markatou, M., Tian, H., Biswas, S., & Hripcsak, G. (2005). Analysis of variance of cross-validation estimators of the generalization error. *Journal of Machine Learning Research*, 6(39), 1127-1168. <https://doi.org/10.7916/D86D5R2X>

- Purnama, A., Lukman, K., Rudiman, R., Prasetyo, D., Fuadah, Y., Nugraha, P., & Candrawinata, V. S. (2023). The prognostic value of COX-2 in predicting metastasis of patients with colorectal cancer: A systematic review and meta-analysis. *Heliyon*, 9(10), Article e21051. <https://doi.org/10.1016/j.heliyon.2023.e21051>
- Raschka, S. (2018). *Model evaluation, model selection, and algorithm selection in machine learning*. arXiv. <https://arxiv.org/abs/1811.12808>
- Shi, C. R., & Adnan, R. (2014). Modified cross-validation as a method for estimating parameter. *AIP Conference Proceedings*, 1635(1), 724-731. <https://doi.org/10.1063/1.4903662>
- Smits, L. J. H., Vink-Börger, E., van Lijnschoten, G., Focke-Snieders, I., van der Post, R. S., Tuynman, J. B., van Grieken, N. C. T., & Nagtegaal, I. D. (2022). Diagnostic variability in the histopathological assessment of advanced colorectal adenomas and early colorectal cancer in a screening population. *Histopathology*, 80(5), 790-798. <https://doi.org/10.1111/his.14601>
- Sukprasert, A. (2023). *Data mining with RapidMiner Studio software* (5th ed.). Mahasarakham University. (in Thai)
- Thitima. (2023, January 13). *Public and private sectors collaborate to develop a health policy lab for value-based healthcare services for cancer patients in Thailand (Value-based Health Care Policy Lab)*. Health Systems Research Institute. <https://wwwold.hsri.or.th/media/news/detail/14222> (in Thai)
- Verma, J., Nath, M., Tripathi, P., & Saini, K. K. (2017). Analysis and identification of kidney stone using Kth nearest neighbour (KNN) and support vector machine (SVM) classification techniques. *Pattern Recognition and Image Analysis*, 27(3), 574-580. <https://doi.org/10.1134/S1054661817030294>
- Vommi, A. M., & Battula, T. K. (2023). A hybrid filter-wrapper feature selection using Fuzzy KNN based on Bonferroni mean for medical datasets classification: A COVID-19 case study. *Expert Systems with Applications*, 218, Article 119612. <https://doi.org/10.1016/j.eswa.2023.119612>
- Wilimitis, D., & Walsh, C. G. (2023). Practical considerations and applied examples of cross-validation for model development and evaluation in health care: Tutorial. *JMIR AI*, 2, Article e49023. <https://doi.org/10.2196/49023>
- Wirth, R., & Hipp, J. (2000). CRISP-DM: Towards a standard process model for data mining. *Proceedings of the 4th international conference on the practical applications of knowledge discovery and data mining* (pp. 29-39).
- World Cancer Research Fund International. (n.d.). *Colorectal cancer statistics*. World

Cancer Research Fund International. <https://www.wcrf.org/preventing-cancer/cancer-statistics/colorectal-cancer-statistics/>

World Health Organization. (2023, July 11). *Colorectal cancer*. World Health Organization. <https://www.who.int/news-room/fact-sheets/detail/colorectal-cancer>

Xu, Y., Jiao, L., Wang, S., Wei, J., Fan, Y., Lai, M., & Chang, E. I. (2013). Multi-label classification for colon cancer using histopathological images. *Microscopy Research and Technique*, 76(12), 1266-1277. <https://doi.org/10.1002/jemt.22294>

Zerouaoui, H., & Idri, A. (2021). Reviewing machine learning and image processing based decision-making systems for breast cancer imaging. *Journal of Medical Systems*, 45(1), Article 8. <https://doi.org/10.1007/s10916-020-01689-1>