

A Model for Screening Patients with Kidney Stones using Image Classification Techniques

Anupong Sukprasert^{1,*}, Namfon Dangphongthong¹ and Theerayuthr Chandra¹

ABSTRACT

This study aims to develop a model for screening kidney stone disease using CT scan images from 6,454 diagnostic exams. We analyzed these images employing the Cross Industry Standard Process for Data Mining (CRISP-DM) and utilized four different image classification techniques to build and assess the model's performance. The accuracy of each method was compared using the Welch Test at a significance level of 0.01. The results indicated significant differences in the mean accuracies across the techniques. Further multiple comparisons using Dunnett's T3 test showed statistically significant differences between each technique pair at the 0.01 significance level. Among the techniques evaluated, k-Nearest Neighbors yielded the highest average performance metrics: an accuracy of 92.64%, an f-measure of 80.72%, a sensitivity of 76.84%, and a specificity of 96.34%. This method is deemed most suitable for developing a preliminary screening model to assess the risk of kidney stones, thereby potentially reducing the diagnostic workload on medical personnel and accelerating the diagnostic process.

Keywords: Image Classification Techniques, Kidney stone screening model, The comparison of model efficiency.

Published Online: 4 June 2024

ISSN: 2730-3829

Anupong Sukprasert^{1,*}

¹Maharakham Business School,
Maharakham University
(anupong.s@acc.msu.ac.th)

Namfon Dangphongthong¹

¹Maharakham Business School,
Maharakham University
(62010912532@msu.ac.th)

Theerayuthr Chandra¹

¹Maharakham Business School,
Maharakham University
(62010912576@msu.ac.th)

^{*}Corresponding Author

Received date: 15 March 2023

Revised date#1: 26 December 2023

Revised date#2: 24 April 2024

Revised date#3: 25 April 2024

Accepted date: 27 May 2024

การสร้างตัวแบบสำหรับการคัดกรองผู้ป่วยโรคหัวใจโดยใช้เทคนิคการ จำแนกประเภทข้อมูลภาพ

อนุพงศ์ สุขประเสริฐ^{1,*} น้ำฝน ดังโพนทอง¹ และธีรยุทธ จันทรา¹

บทคัดย่อ

งานวิจัยนี้มีวัตถุประสงค์เพื่อสร้างตัวแบบสำหรับการคัดกรองผู้ป่วยโรคหัวใจใน
ไต ด้วยภาพถ่าย CT Scan ของผู้ที่เข้ารับการวินิจฉัย จำนวน 6,454 ภาพ โดยนำมา
วิเคราะห์ตามกระบวนการมาตรฐานการทำเหมืองข้อมูล และใช้เทคนิคการจำแนก
ประเภทข้อมูลภาพ 4 เทคนิค สำหรับการสร้างตัวแบบและวิเคราะห์ประสิทธิภาพ
ของตัวแบบในการจำแนกประเภทข้อมูลภาพ จากนั้นใช้ค่าความแม่นยำมาทดสอบ
ความแตกต่างของเทคนิคการจำแนกประเภทข้อมูลภาพด้วยสถิติ Welch Test ที่
ระดับนัยสำคัญ 0.01 ผลการทดสอบพบว่า ค่าเฉลี่ยของความแม่นยำที่ได้จากเทคนิค
การจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค มีความแตกต่างกัน ผู้วิจัยจึงได้ทำการ
ทดสอบแบบพหุคูณด้วยวิธี Dunnett's T3 พบว่าค่าเฉลี่ยของความแม่นยำของเทคนิค
แต่ละคู่มีความแตกต่างกัน อย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 และเมื่อนำค่า
ประสิทธิภาพของการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิคมาเปรียบเทียบกัน
พบว่าเทคนิคที่ให้ค่าเฉลี่ยของประสิทธิภาพการจำแนกข้อมูลสูงที่สุดคือ เทคนิค
เพื่อนบ้านใกล้ที่สุด ซึ่งมีค่าความแม่นยำเท่ากับ 92.64% ค่าประสิทธิภาพโดยรวม
เท่ากับ 80.72% ค่าความไวเท่ากับ 76.84% และค่าจำเพาะเท่ากับ 96.34% ซึ่งเป็น
เทคนิคที่มีความเหมาะสมสำหรับนำไปสร้างตัวแบบเพื่อคัดกรองความเสี่ยงผู้ป่วยที่มี
โอกาสเป็นโรคหัวใจในไตเบื้องต้น เพื่อลดภาระให้กับบุคลากรทางแพทย์ในการตรวจ
คัดกรอง ทำให้การวินิจฉัยมีความรวดเร็วยิ่งขึ้น

คำสำคัญ: เทคนิคการจำแนกข้อมูลด้วยภาพ ตัวแบบการคัดกรองโรคหัวใจในไต
การเปรียบเทียบประสิทธิภาพของตัวแบบ

Published Online: 4 มิถุนายน 2567

ISSN: 2730-3829

อนุพงศ์ สุขประเสริฐ^{1,*}

¹ คณะการบัญชีและการจัดการ

มหาวิทยาลัยมหาสารคาม

(anupong.s@acc.msu.ac.th)

น้ำฝน ดังโพนทอง¹

¹ คณะการบัญชีและการจัดการ

มหาวิทยาลัยมหาสารคาม

(62010912532@msu.ac.th)

ธีรยุทธ จันทรา¹

¹ คณะการบัญชีและการจัดการ

มหาวิทยาลัยมหาสารคาม

(62010912576@msu.ac.th)

**Corresponding Author*

Received date: 15 มีนาคม 2566

Revised date: 26 ธันวาคม 2566

Revised date: 24 เมษายน 2567

Revised date: 25 เมษายน 2567

Accepted date: 27 พฤษภาคม 2567

1. บทนำ

โรคนี้ในไตเป็นโรคที่พบได้บ่อยในทุกเพศทุกวัยทุกประเทศทั่วโลก ความชุกของโรคนี้ในไตทั่วโลกพบประมาณร้อยละ 2 – 20 ซึ่งพบมากในเพศชายมากกว่าเพศหญิง และพบมากในช่วงอายุ 30 - 40 ปี ของผู้ป่วยในประเทศแถบตะวันตกพบอุบัติการณ์การเป็นโรคนี้ในไต ประมาณร้อยละ 5 - 10 (ยุพิน ปักกะสังข์ นงลักษณ์ เมธากาญจนศักดิ์ และอัมพรพรรณ ธีรานูตร, 2562) และมีแนวโน้มเพิ่มขึ้นแบบทวีคูณทั่วโลก จากการเปลี่ยนแปลงสภาพแวดล้อมจากภาวะโลกร้อน สังคมอุตสาหกรรม การระบาดของโรคอ้วนและโรคเมตาบอลิกซินโดรม โดยคาดการณ์ว่า ในปี ค.ศ. 2050 อัตราชุกของโรคนี้ในไตจะสูงถึงร้อยละ 56 และเพิ่มขึ้นเป็นร้อยละ 70 ในปี ค.ศ. 2095 สำหรับประเทศในเขตร้อนพบได้ถึงร้อยละ 20 แม้ในประเทศที่พัฒนาแล้วอย่างสหรัฐอเมริกา พบสูงถึงร้อยละ 20 ส่วนในทวีปเอเชีย เช่น ประเทศซาอุดีอาระเบียพบสูงถึงร้อยละ 21 สำหรับประเทศไทยพบความชุกโรคนี้ในไตในภาคตะวันออกเฉียงเหนือสูงกว่าภาคอื่น ๆ ประมาณร้อยละ 15 - 16.9 (ลักขณา พรหมกสิกร, 2558) สำหรับในประเทศไทยอุบัติการณ์การเป็นโรคนี้ในไตจะพบสูงมากในแถบภาคตะวันออกเฉียงเหนือ คิดเป็นร้อยละ 15 - 16.9 โดยอัตราการเกิดนี้ 38 คนต่อประชากร 10,000 คน และขนาดของก้อนนี้อาจมีขนาดที่แตกต่างกัน อาจมีเพียงก้อนเดียวหรือหลายก้อนก็ได้ โดยก้อนนี้ที่เกิดขึ้นในไตประกอบด้วยหินปูนแคลเซียม กับสารเคมีอื่น เช่น ออกซาเลต กรดยูริก เป็นต้น ดังนั้นการเกิดนี้จึงมีส่วนเกี่ยวข้องกับภาวะที่มีแคลเซียมในปัสสาวะมากผิดปกติ ทำให้เกิดภาวะแทรกซ้อนของโรคนี้ในไตหลายอย่าง อาทิ การติดเชื้อในระบบทางเดินปัสสาวะ (Urinary Tract Infection: UTI) กรวยไตอักเสบ (Pyelonephritis) ภาวะไตบวมน้ำ (Hydronephrosis: HDN) ไตเรื้อรังระยะสุดท้าย (End Stage Renal Disease: ESRD) (ยุพิน ปักกะสังข์ นงลักษณ์ เมธากาญจนศักดิ์ และอัมพรพรรณ ธีรานูตร, 2562) โรคนี้ในไตสามารถเกิดซ้ำได้แม้จะได้รับการผ่าตัดรักษาแล้วก็ตาม เนื่องจากสาเหตุของนี้ไตเกี่ยวข้องกับพฤติกรรมกรับประทานอาหาร การติดเชื้อทางเดินของปัสสาวะ หรือความผิดปกติของไตโดยกำเนิด และอาจพบร่วมกับคนที่มีกรดยูริกในเลือดสูง โรคเก๊าท์ และต่อมพาราไทรอยด์ทำงานมากเกินไป

จากปัญหาข้างต้นผู้วิจัยจึงมีแนวคิดนำทฤษฎีการประมวลผลภาพมาวิเคราะห์ข้อมูลจากภาพการตรวจวินิจฉัยโรคด้วยเครื่องเอกซเรย์คอมพิวเตอร์ (Computerized Tomography Scan: CT scan) เพื่อคัดกรองผู้ป่วยโรคนี้ในไต โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพ 4 เทคนิค ได้แก่ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machines: SVM) เทคนิคเพื่อนบ้านใกล้ที่สุด (k-Nearest Neighbors: k-NN) เทคนิคต้นไม้ป่าสุ่ม (Random Forest) และเทคนิคนาอิวเบย์ (Naive Bayes) จากนั้นทำการเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในไตเพื่อคัดเลือกตัวแบบที่มีประสิทธิภาพการจำแนกประเภทข้อมูลที่ดีที่สุด เพื่อนำมาสร้างเป็นตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในไตเบื้องต้นจากภาพถ่าย CT scan และหาแนวทางในการรักษาโรคได้ตรงกับอาการมากที่สุด อีกทั้งยังเป็นการช่วยสนับสนุนการตัดสินใจให้กับแพทย์ในการวินิจฉัยโรคได้อย่างรวดเร็วและมีความแม่นยำที่สุด

วัตถุประสงค์

1. เพื่อสร้างตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในไตด้วยภาพ CT scan
2. เพื่อเปรียบเทียบประสิทธิภาพของตัวแบบที่ใช้สำหรับการคัดกรองผู้ป่วยโรคนี้ในไตด้วยภาพ CT scan

2. วิธีการวิจัย

การดำเนินงานวิจัยสำหรับการสร้างตัวแบบและการเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในไตด้วยภาพ CT scan โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพ ผู้วิจัยได้ดำเนินการโดยอ้างอิงตามกระบวนการมาตรฐานในการทำเหมืองข้อมูล (CRISP-DM) (อนุพงศ์ สุขประเสริฐ, 2564) ทั้ง 6 ขั้นตอน ดังนี้

1) การทำความเข้าใจเกี่ยวกับธุรกิจ (Business Understanding)

โรคไตในไตเป็นปัญหาสาธารณสุขที่พบบ่อยและพบมากทั่วโลก รวมทั้งประเทศไทย สาเหตุการเกิดโรคเกิดจากการอัมตัวของสารเมตาบอลิซึมและของเสียในปัสสาวะ ที่มาจากการบริโภคอาหารที่เป็นปัจจัยเสี่ยงต่าง ๆ ทำให้เกิดการก่อตัวเป็นตะกอนขนาดเล็กก่อนรวมกันเป็นก้อนผลึกแข็ง การเกิดก้อนนี้ไว้จะทำให้เกิดการอุดตัน ขัดขวาง ทำให้มีการคั่งและการไหลย้อนกลับของปัสสาวะ และเกิดภาวะแทรกซ้อนตามมาคือ การติดเชื้อในระบบทางเดินปัสสาวะ เป็นสาเหตุสำคัญให้เกิดการติดเชื้อเข้าสู่กระแสเลือด (Septicemia) ทำให้ผู้ป่วยเสียชีวิตได้ โดยเฉพาะในผู้สูงอายุและพบว่าเป็นสาเหตุอันดับ 4 ของโรคไตเรื้อรังในประเทศไทย ซึ่งการมีนิ่วไตตั้งแต่ 1 ก้อนขึ้นไป จะเพิ่มโอกาสเกิดโรคไตเรื้อรังระยะสุดท้าย (End Stage Renal Disease: ESRD) มากกว่าคนที่ไม่มีนิ่วถึง 2.16 เท่า แม้โรคไตจะเป็นสาเหตุสำคัญของการเกิดโรคภัยร้ายแรงต่าง ๆ แต่เป็นโรคที่สามารถป้องกัน ยับยั้ง และรักษาได้ด้วยการปรับเปลี่ยนพฤติกรรมสุขภาพที่เหมาะสม (ลักขณา พรหมกลีกร, 2558) ผู้วิจัยจึงมีแนวคิดในการพัฒนาตัวแบบสำหรับคัดกรองผู้ป่วยโรคไตในไต โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพจากภาพถ่าย CT scan ซึ่งเป็นวิธีการอีกทางเลือกที่จะสามารถวินิจฉัยโรคได้อย่างรวดเร็วและลดภาระให้กับบุคลากรทางการแพทย์ อีกทั้งยังสามารถนำผลการวิเคราะห์ที่ได้ไปสนับสนุนการตัดสินใจของแพทย์ในการวินิจฉัยโรคได้อย่างรวดเร็วมากยิ่งขึ้น

2) การทำความเข้าใจเกี่ยวกับข้อมูล (Data Understanding)

งานวิจัยนี้ใช้ชุดข้อมูลจริงเกี่ยวกับภาพของผู้ป่วยโรคไตในไตที่ถูกบันทึกไว้ในปี พ.ศ.2565 ซึ่งมีจำนวนทั้งสิ้น 6,454 ภาพ โดยแบ่งเป็น 2 โฟลเดอร์ ได้แก่ โฟลเดอร์ S สำหรับการจัดเก็บข้อมูลภาพของผู้ป่วยโรคไตในไต จำนวน 1,377 ภาพ และโฟลเดอร์ N สำหรับการจัดเก็บข้อมูลภาพของคนปกติ จำนวน 5,077 ภาพ และถูกจัดเก็บในรูปแบบไฟล์ภาพนามสกุล .jpg แหล่งข้อมูลได้มาจากเว็บไซต์ <https://www.kaggle.com> ที่ได้ถูกรวบรวมโดย Mehedi Hasan นักเทคโนโลยีการแพทย์ ซึ่งรวบรวมจากระบบเก็บถาวรรูปภาพและการสื่อสาร (Picture Archiving and Communication System: PACS) จากโรงพยาบาลต่าง ๆ ในกรุงธากา ประเทศบังกลาเทศ (Md Humaion Kabir Mehedi, 2022)

3) การเตรียมข้อมูล (Data Preparation)

ผู้วิจัยได้นำเข้าข้อมูลภาพมาทำการเตรียมข้อมูลให้อยู่ในรูปแบบข้อมูลที่มีโครงสร้างก่อนการวิเคราะห์ด้วยโปรแกรม RapidMiner Studio ซึ่งเป็นโปรแกรมที่ออกแบบมาสำหรับการวิเคราะห์ข้อมูล การทำเหมืองข้อมูล (Data mining) การทำเหมืองข้อความ (Text mining) และการทำเหมืองข้อมูลภาพ (Image mining) ซึ่งเป็นกระบวนการ (Process) ที่กระทำกับข้อมูลขนาดใหญ่เพื่อค้นหารูปแบบ แนวทาง และความสัมพันธ์ที่ซ่อนอยู่ในชุดข้อมูลนั้นโดยอาศัยหลักสถิติ การรู้จำ การเรียนรู้ของเครื่อง และหลักคณิตศาสตร์เพื่อให้ได้สารสนเทศออกมา (อนุพงศ์ สุขประเสริฐ, 2564) โดยผู้วิจัยได้ทำการแบ่งเก็บภาพแต่ละกลุ่มไว้ในแต่ละโฟลเดอร์ 2 โฟลเดอร์ย่อย ได้แก่ โฟลเดอร์ S สำหรับการจัดเก็บข้อมูลภาพของผู้ป่วยโรคไตในไต และโฟลเดอร์ N สำหรับการจัดเก็บข้อมูลภาพของคนที่มีไตปกติ โดยมีความละเอียดภาพอยู่ที่ 512x512 พิกเซล ทั้ง 2 โฟลเดอร์ ดัง Figure 1

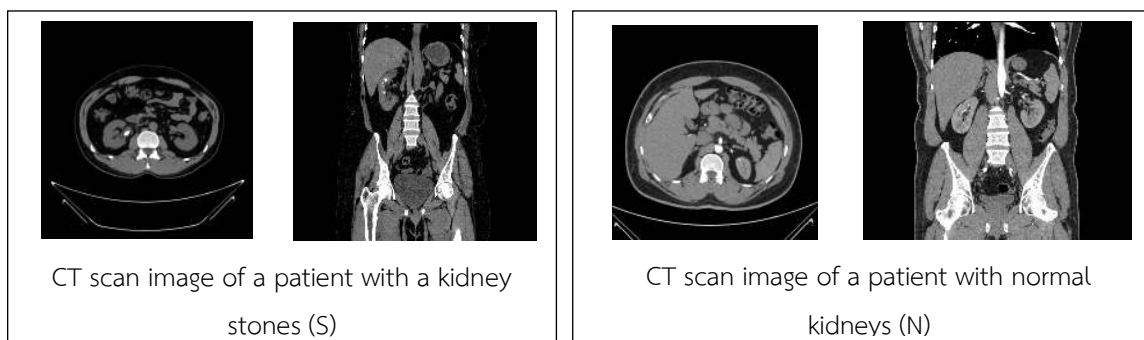


Figure 1 Examples of CT scan images taken from CT kidney dataset

จาก Figure 1 แสดงภาพถ่าย CT scan ของผู้ที่เข้ารับการวินิจฉัยคัดกรองโรคนิ่วในไต ซึ่งแบ่งออกเป็น 2 โพลเดอร์ คือ โพลเดอร์ S สำหรับการจัดเก็บข้อมูลภาพของผู้ป่วยโรคนิ่วในไต และโพลเดอร์ N สำหรับการจัดเก็บข้อมูลภาพของคนที่มีไตปกติ

จากนั้นแปลงข้อมูลให้อยู่ในรูปแบบข้อมูลที่มีโครงสร้าง (Structured Data) สำหรับงานวิจัยนี้ผู้วิจัยได้ใช้การแยกคุณลักษณะโดยธรรมชาติสำหรับการทำเหมืองรูปภาพในระดับสากล (Global Level) เพื่อแยกคุณสมบัติส่วนกลางออกจากภาพถ่าย CT scan และบ่งบอกถึงความแตกต่างของภาพ ซึ่งเหมาะสำหรับการจำแนกประเภทของภาพและการกำหนดคุณสมบัติของภาพโดยรวม การคำนวณค่าต่าง ๆ จากรูปภาพแบบโหมดระดับสีเทา (Gray scale) ซึ่งในงานวิจัยนี้จะทำการแยกคุณสมบัติของภาพทั้งหมดโดยใช้ค่าคุณสมบัติทั้ง 7 ค่า ได้แก่ ค่าสีเทาสูงสุด (Max Gray Value) ค่าสีเทาขั้นต่ำ (Min Gray Value) ค่าเฉลี่ย (Mean) ค่ามัธยฐาน (Median) ค่ามาตรฐานสำหรับแกน X (Normalized X) ค่ามาตรฐานสำหรับแกน Y (Normalized Y) และค่าส่วนเบี่ยงเบนมาตรฐาน (Standard Deviation) โดยใช้โปรแกรม RapidMiner Studio version 10 ที่มีการติดตั้งส่วนขยายการทำเหมืองรูปภาพ (Image Mining Extension for RapidMiner: IMMI) สำหรับทำการแปลงข้อมูลรูปภาพให้อยู่ในรูปแบบของข้อมูลที่มีโครงสร้างโดยการแยกคุณลักษณะ ดังแสดงใน Table 1

Table 1 A example of images data transformed to a structured data format

Row No.	Label	Max Gray Value Global statistics	Min Gray Value Global statistics	Mean Global statistics	Median Global statistics	Normalized X Center of Mass Global statistics	Normalized Y Center of Mass Global statistics	Standard Deviation Global statistics
1	Normal (N)	255	0	50.504	15	0.476	0.559	63.187
2	Normal (N)	255	0	49.843	10	0.476	0.559	62.540
3	Normal (N)	255	0	42.744	0	0.484	0.568	59.650
4	Kidney Stones (S)	255	0	46.686	0	0.495	0.556	65.225
5	Kidney Stones (S)	255	0	46.543	0	0.495	0.555	64.961
:	:	:	:	:	:	:	:	:

4) การสร้างแบบจำลอง (Modeling)

เมื่อข้อมูลรูปภาพถูกแปลงให้อยู่ในรูปแบบข้อมูลที่มีโครงสร้างแล้ว ขั้นตอนต่อมา คือการสร้างตัวแบบสำหรับการสร้างตัวแบบการคัดกรองผู้ป่วยโรคนิ่วในไต โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพ (Image Classification Technique) จำนวน 4 เทคนิค ประกอบด้วย เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machines: SVM) เทคนิคเพื่อนบ้านใกล้ที่สุด (k-Nearest Neighbors: k-NN) เทคนิคนาอิวเบย์ (Naive Bayes) และเทคนิคต้นไม้ป่าสุ่ม (Random Forest) มีรายละเอียด ดังนี้

4.1 เทคนิคซัพพอร์ตเวกเตอร์แมชชีน (Support Vector Machines: SVM) เป็นวิธีที่ใช้ในการแบ่งคุณลักษณะของข้อมูลออกเป็น 2 กลุ่มขึ้นไป โดยการสร้างเส้นแบ่ง (Hyperplane) ซึ่งเป็นเส้นตรง เพื่อระบุเส้นแบ่งที่ดีที่สุดระหว่าง 2 กลุ่มข้อมูล โดยการเพิ่มเส้นขอบ (Margin) ออกไปทั้งสองข้างจนสัมผัสกับข้อมูลของกลุ่มตัวอย่างที่ใกล้ที่สุด การปรับค่าสมการเพื่อค้นหาเส้นแบ่งที่เหมาะสมที่สุดใช้วิธีของเคอร์เนล (Kernel) ที่เป็นฟังก์ชันใช้ในการแปลงข้อมูลเชิงเส้น (Linear) ให้เป็นรูปแบบที่ซับซ้อนเพื่อสามารถแยกแยะข้อมูลที่ไม่เชิงเส้นได้อย่างมีประสิทธิภาพ (นพรัตน์ นนทศิริ, พิศณุชัย จิตวณิชกุล และกริช สมกันธา, 2565)

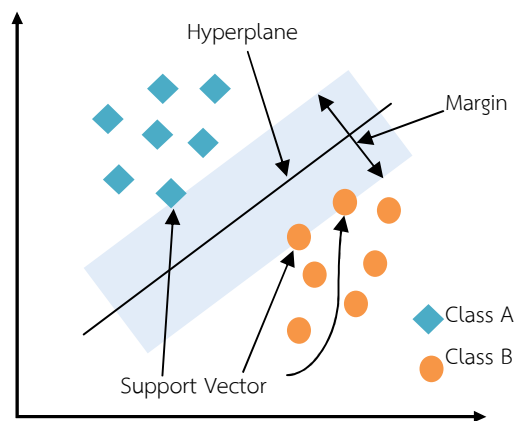


Figure 2 Support Vector Machines Algorithm

4.2 เทคนิคเพื่อนบ้านใกล้ที่สุด (k-Nearest Neighbors : k-NN) เป็นเทคนิคที่ใช้ในการจัดแบ่งคลาสของข้อมูล โดยการเลือกคลาสที่จะแทนเงื่อนไขหรือกรณีใหม่ ๆ โดยการพิจารณาจากจำนวนบางจำนวนของกรณีหรือเงื่อนไขที่มีความเหมือนหรือใกล้เคียงกันมากที่สุด วิธีการนี้จะนับจำนวนเงื่อนไขและกำหนดค่า k ให้เป็นเลขคี่ แล้วคำนวณระยะห่างระหว่างข้อมูลที่ต้องการพิจารณากับกลุ่มข้อมูลตัวอย่าง เรียงลำดับระยะห่างและเลือกชุดข้อมูลที่ใกล้จุดที่ต้องการพิจารณาตามจำนวน k ที่กำหนด จากนั้นจะพิจารณาข้อมูล k ชุด และกำหนดคลาสให้กับจุดที่พิจารณาซึ่งอยู่ใกล้จุดพิจารณามากที่สุด (อนุชิต ละอองคำ และพิศณุ คุ้มชัย, 2563)

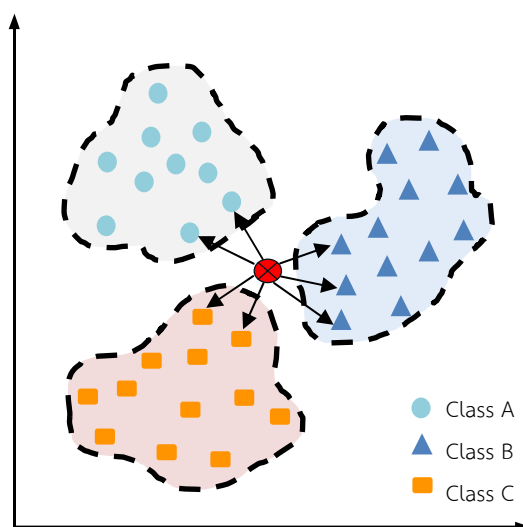


Figure 3 k-Nearest Neighbors Algorithm

4.3 เทคนิคนาอิวเบย์ (Naive Bayes) เป็นการจำแนกประเภทข้อมูลที่เหมาะสมกับกรณีของเซตตัวอย่างที่มีจำนวนมากและคุณสมบัติของตัวอย่างมีความเป็นอิสระต่อกัน โดยอาศัยหลักการของทฤษฎีเบย์เพื่อคำนวณหาความน่าจะเป็น (Probability) ที่ตัวอย่างในข้อมูลของเซตตัวอย่าง ซึ่งหาได้จากสมการที่ (1) และ (2) (นพรัตน์ นนท์ศิริ และคณะ, 2565)

$$P(c|x) = \frac{P(x|c)P(c)}{P(x)} \quad (1)$$

$$\text{และ} \quad P(c|x) = P(x_1|c) \times P(x_2|c) \times \dots \times P(x_n|c) \times P(c) \quad (2)$$

โดยที่ c คือ คลาสของข้อมูล
 x คือ แอตทริบิวต์
 P คือ ความน่าจะเป็นของข้อมูล
 $P(c|x)$ คือ ความน่าจะเป็นที่ข้อมูลที่มีแอตทริบิวต์เป็น x จะมีคลาส c
 $P(x|c)$ คือ ความน่าจะเป็นที่ข้อมูลที่มีคลาส c และมีแอตทริบิวต์ x
 $P(c)$ คือ จำนวนคลาสที่อาจจะเกิดขึ้นหารด้วยจำนวนคลาสทั้งหมดของคลาส c
 $P(x)$ คือ จำนวนแอตทริบิวต์ทั้งหมด

4.4 เทคนิคต้นไม้ป่าสุ่ม (Random Forest) เป็นการสร้างแบบจำลองทางสถิติและเป็นโมเดลอีกประเภทหนึ่งของการเรียนรู้ของเครื่องจักร (Machine Learning) ซึ่งนำมาใช้กับข้อมูลที่มีหลายตัวแปร (Multivariate) หรือข้อมูลที่มีความซับซ้อน พัฒนาขึ้นจากต้นไม้ตัดสินใจใช้หลักการของการแบ่งกลุ่ม โดยการรวมกันของข้อมูลหลาย ๆ ต้นไม้ตัดสินใจ (Qadri, 2021)

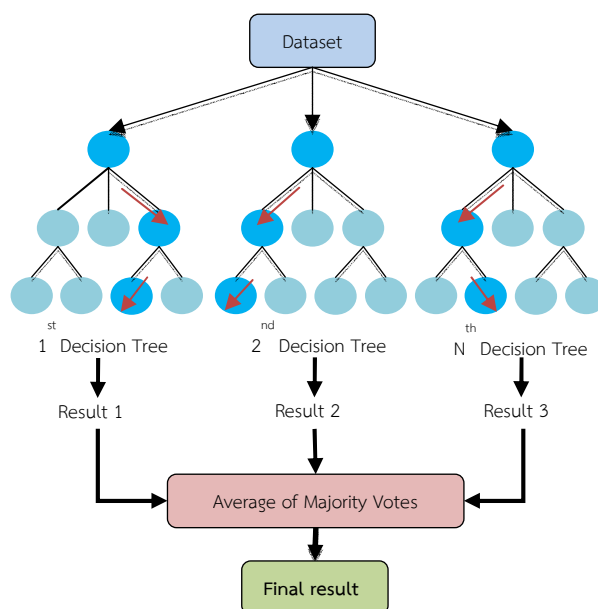


Figure 4 Random Forest Algorithm

5) การประเมินผล (Evaluation)

การประเมินประสิทธิภาพของตัวแบบผู้วิจัยได้ทำการแบ่งข้อมูลออกเป็น 2 ส่วน คือ ข้อมูลส่วนแรกใช้สำหรับการสร้างแบบจำลอง (Training Set) และข้อมูลอีกส่วนใช้สำหรับการทดสอบประสิทธิภาพของตัวแบบ (Testing Set) ด้วยวิธี Cross validation โดยแบ่งข้อมูลออกเป็น 10 ส่วน เท่า ๆ กัน (10-folds cross validation) และทำการทดสอบประสิทธิภาพของการจำแนกประเภทข้อมูลภาพด้วยค่าความแม่นยำ (Accuracy) ค่าประสิทธิภาพโดยรวม (F-measure) ค่าความไว (Sensitivity) และค่าจำเพาะ (Specificity) โดยสร้างเป็นเมทริกซ์ความสับสน (Confusion Matrix) เพื่อการเก็บข้อมูลที่เกี่ยวข้องกับการแบ่งแยกข้อมูลจริงกับข้อมูลที่เกิดจากการทำนายด้วยระบบการแบ่งแยก (Classification Systems) ดัง Table 2 ซึ่งเมทริกซ์ความสับสนเป็นตารางที่แสดงผลลัพธ์ของตัวจำแนกประเภทข้อมูล (Classifier) ว่าทำนายผลถูกต้องหรือไม่ โดยเปรียบเทียบผลลัพธ์ที่คาดการณ์ (Predicted) กับผลลัพธ์ที่เป็นจริง (Actual)

Table 2 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)

5.1 ค่าความแม่นยำ (Accuracy) ซึ่งเป็นค่าที่ได้จากวิธีการทดสอบเพื่อหาค่าพยากรณ์ความถูกต้องของข้อมูลโดยคิดเป็นร้อยละ (%) ใช้สมการที่ (3) เพื่อเปรียบเทียบประสิทธิภาพของตัวแบบ โดยผู้วิจัยจะนำตัวแบบที่มีความเหมาะสมที่สุดสำหรับนำมาสร้างตัวแบบการคัดกรองผู้ป่วยโรคนี้ในต่อไป

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}} \quad (3)$$

5.2 ค่าประสิทธิภาพโดยรวม (F-measure) คือ ค่าที่เกิดจากการเปรียบเทียบระหว่าง ค่า Precision และ ค่า Recall ของแต่ละคลาสเป้าหมาย ดังสมการที่ (4) และ (5)

$$\text{F-measure YES} = \frac{(2 * \text{Precision(YES)} * \text{Recall(YES)})}{(\text{Precision(YES)} * \text{Recall(YES)})} \quad (4)$$

$$\text{F-measure NO} = \frac{(2 * \text{Precision(NO)} * \text{Recall(NO)})}{(\text{Precision(NO)} * \text{Recall(NO)})} \quad (5)$$

5.3 ค่าความไว (Sensitivity) คือ ค่าที่แบบจำลองสามารถพยากรณ์ข้อมูลผู้ป่วยที่เกิดโรคได้อย่างถูกต้องต่อผู้ป่วยที่เกิดโรคจริง ดังสมการที่ (6)

$$\text{Sensitivity} = \frac{\text{TP}}{\text{TP} + \text{FN}} \quad (6)$$

5.4 ค่าจำเพาะ (Specificity) คือ ค่าที่แบบจำลองสามารถพยากรณ์ข้อมูลผู้ป่วยที่ไม่เกิดโรคได้อย่างถูกต้องต่อผู้ป่วยที่ไม่เกิดโรคจริง ดังสมการที่ (7)

$$\text{Specificity} = \frac{\text{TN}}{\text{TN} + \text{FP}} \quad (7)$$

โดยที่ True Positive (TP) คือ สิ่งที่ทำนายตรงกับสิ่งที่เกิดขึ้นจริงในกรณีทำนายว่าจริง และสิ่งที่เกิดขึ้นคือจริง
 True Negative (TN) คือ สิ่งที่ทำนายตรงกับสิ่งที่เกิดขึ้นจริงในกรณีทำนายว่าไม่จริง และสิ่งที่เกิดขึ้นคือไม่จริง
 False Positive (FP) คือ สิ่งที่ทำนายไม่ตรงกับสิ่งที่เกิดขึ้นจริงคือทำนายว่าจริง แต่สิ่งที่เกิดขึ้นคือไม่จริง
 False Negative (FN) คือ สิ่งที่ทำนายไม่ตรงกับสิ่งที่เกิดขึ้นจริงคือทำนายว่าไม่จริง แต่สิ่งที่เกิดขึ้นคือจริง

6) การนำแบบจำลองไปใช้งาน (Deployment)

เมื่อได้ทำการวิเคราะห์ครบ 5 ขั้นตอนเสร็จเรียบร้อยแล้ว จะได้ผลลัพธ์ของเทคนิคที่เหมาะสมสำหรับการสร้างตัวแบบการพยากรณ์ในการคัดกรองผู้ป่วยโรคนี้ในไตด้วยด้วยรูปภาพ ซึ่งตัวแบบที่ได้จะสามารถนำไปใช้สำหรับสนับสนุนการตัดสินใจสำหรับแพทย์ในการคัดกรองผู้ป่วยโรคนี้ในไต ทำให้การวินิจฉัยมีความแม่นยำและรวดเร็วยิ่งขึ้น เพื่อสนับสนุนนโยบายด้านการวางแผนการรักษา และป้องกันควบคุมความเสี่ยงของการเป็นโรคนี้ในไตได้นอกจากนี้

ยังสามารถนำไปพัฒนาต่อยอดเป็นระบบสารสนเทศสำหรับการคัดกรองผู้ที่มีความเสี่ยงต่อการเป็นโรคนี้ในโต และเตรียมความพร้อมสำหรับทีมสหวิชาชีพให้เพียงพอต่อจำนวนผู้ป่วยที่จะเกิดขึ้นในอนาคต

3. ผลการวิจัย

ผู้วิจัยได้ทำการวิเคราะห์ข้อมูลเพื่อสร้างตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในโตด้วยภาพถ่าย CT Scan โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค ได้แก่ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน เทคนิคเพื่อนบ้านใกล้ที่สุด เทคนิคนาอ์ฟเบย์ และเทคนิคต้นไม้ป่าสุ่ม ซึ่งได้ทำการรันข้อมูลสำหรับสร้างตัวแบบทั้งหมด 30 ครั้ง และทำการปรับแต่งไฮเปอร์พารามิเตอร์ (Hyperparameter Tuning) ซึ่งเป็นค่าปกติที่ทางโปรแกรม RapidMiner Studio ได้กำหนดไว้ให้ในแต่ละเทคนิค ดังนี้ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน ใช้ฟังก์ชันเคอร์เนลแบบดอท โดยกำหนดค่าพารามิเตอร์ C และ epsilon ให้มีค่าเท่ากับ 0.0 ส่วนค่า convergence epsilon ให้มีค่าเท่ากับ 0.001 และกำหนดให้มีการวนซ้ำสูงสุด (Max iterations) เท่ากับ 100,000 รอบ สำหรับเทคนิคเพื่อนบ้านใกล้ที่สุด มีการกำหนดค่าพารามิเตอร์ k ให้มีค่าเท่ากับ 5 และเทคนิคต้นไม้ป่าสุ่มมีการกำหนดค่าพารามิเตอร์ของค่าความลึกของต้นไม้แต่ละต้น (Maximal depth) ให้มีค่าเท่ากับ 10 และจำนวนต้นไม้สูงสุด (number of trees) อยู่ที่ 1,000 ต้น จากนั้นทดสอบประสิทธิภาพของตัวแบบด้วยค่าความแม่นยำ ค่าประสิทธิภาพโดยรวม ค่าความไว และค่าจำเพาะ และมีการวิเคราะห์สำหรับทดสอบความแปรปรวน (Analysis of Variance: ANOVA) เพื่อทดสอบสมมติฐานว่าเทคนิคการจำแนกประเภทข้อมูลภาพที่ใช้ในงานวิจัยนี้มีประสิทธิภาพของการจำแนกประเภทข้อมูลแตกต่างกันหรือไม่ โดยใช้ค่าเฉลี่ยของค่าความแม่นยำมาทำการทดสอบ ที่ระดับนัยสำคัญทางสถิติ 0.01 ซึ่งให้ผลลัพธ์ดัง Table 3

Table 3 Analysis of Variance of Image Classification Techniques using Welch Test

Levene Statistic	Sig.	Welch test	Sig.
37.27	0.000*	528787.57	0.000*

* at 0.01 significant level

จาก Table 3 แสดงผลการวิเคราะห์ความเท่ากันของความแปรปรวนของเทคนิคการจำแนกประเภทข้อมูลภาพ 4 เทคนิค ด้วยค่า Levene Statistic ซึ่งมีค่าเท่ากับ 37.27 และค่า Sig เท่ากับ 0.000 (Sig. < 0.01) ซึ่งมีความน้อยกว่าค่าระดับนัยสำคัญที่กำหนด หมายความว่า ความแปรปรวนของค่าเฉลี่ยของความแม่นยำในแต่ละเทคนิคนั้นแตกต่างกัน ที่ระดับนัยสำคัญทางสถิติ 0.01 จึงใช้สถิติ Welch Test ในการทดสอบความแตกต่างของค่าเฉลี่ยของความแม่นยำที่ได้จากเทคนิคการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค พบว่ามีค่า Welch Test มีค่าเท่ากับ 528787.57 และค่า Sig. เท่ากับ 0.000 (Sig. < 0.01) ซึ่งมีความน้อยกว่าค่าระดับนัยสำคัญที่กำหนด หมายความว่าค่าเฉลี่ยของความแม่นยำทั้ง 4 เทคนิคมีความแตกต่างกัน ที่ระดับนัยสำคัญทางสถิติ 0.01 จากนั้นผู้วิจัยได้ทำการทดสอบความแตกต่างของเทคนิคการจำแนกประเภทข้อมูลภาพแบบพหุคูณด้วยวิธี Dunnett T3 ซึ่งให้ผลลัพธ์ดัง Table 4

Table 4 Testing differences in multiple pairwise comparisons of image classification techniques using Dunnett's T3 method.

(I) Classification Techniques	(J) Classification Techniques	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
k-NN	SVM	9.36367	0.15421	0.000*	8.8019	9.9254
	Random Forest	7.16167	0.15360	0.000*	6.6012	7.7221
	Naive Bayes	16.76500	0.15355	0.000*	16.2047	17.3253
SVM	Random Forest	-2.20200	0.01566	0.000*	-2.2579	-2.1461
	Naive Bayes	7.40133	0.01518	0.000*	7.3466	7.4560
Random Forest	Naive Bayes	9.60333	0.00665	0.000*	9.5802	9.6265

* at 0.01 significant level

จาก Table 4 แสดงผลการเปรียบเทียบประสิทธิภาพการจำแนกประเภทข้อมูลภาพแบบพหุคูณด้วยวิธี Dunnett's T3 พบว่าค่าเฉลี่ยของค่าความแม่นยำที่ได้จากการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค ทุกคู่มีความแตกต่างกันอย่างมีนัยสำคัญทางสถิติที่ระดับ 0.01 จากนั้นผู้วิจัยได้วิเคราะห์หาค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของค่าประสิทธิภาพของการจำแนกประเภทข้อมูลภาพของตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในไตด้วยด้วยค่าความแม่นยำ ค่าประสิทธิภาพโดยรวม ค่าความไว และค่าจำเพาะ ดัง Table 5

Table 5 Mean and standard deviation of performance in image classification across all four techniques

Classification Performance	Image Classification Techniques							
	k-Nearest Neighbors		Support Vector Machines		Random Forest		Naive Bayes	
	$\bar{X}$	S.D.	$\bar{X}$	S.D.	$\bar{X}$	S.D.	$\bar{X}$	S.D.
Accuracy	92.64	0.84	83.28	0.080	85.48	0.029	75.88	0.02
F-measure	80.72	0.93	40.08	0.37	48.64	0.13	46.58	0.82
Sensitivity	76.84	1.52	26.30	0.25	32.39	0.11	48.54	0.29
Specificity	96.34	0.10	98.74	0.049	99.88	0.01	83.30	0.06

จาก Table 5 แสดงค่าเฉลี่ยและส่วนเบี่ยงเบนมาตรฐานของประสิทธิภาพการจำแนกประเภทข้อมูลภาพทั้ง 4 เทคนิค พบว่า เทคนิคที่ให้ค่าเฉลี่ยความแม่นยำสูงสุดคือ เทคนิคเพื่อนบ้านใกล้ที่สุด ซึ่งมีค่าเท่ากับ 92.64% รองลงมาคือ เทคนิคต้นไม้ป่าสุ่ม มีเท่ากับ 85.48% และเทคนิคที่ให้ค่าเฉลี่ยความแม่นยำน้อยที่สุด คือ เทคนิคนาอิวเบย์ ซึ่งมีค่าเท่ากับ 75.88% เทคนิคที่ให้ค่าเฉลี่ยค่าประสิทธิภาพโดยรวมสูงสุดคือ เทคนิคเพื่อนบ้านใกล้ที่สุด ซึ่งมีค่าเท่ากับ 80.72% รองลงมาคือ เทคนิคต้นไม้ป่าสุ่ม มีเท่ากับ 48.64% และเทคนิคที่ให้ค่าเฉลี่ยค่าประสิทธิภาพโดยรวมน้อยที่สุดคือ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน ซึ่งมีค่าเท่ากับ 40.08% เทคนิคที่ให้ค่าเฉลี่ยความไวสูงสุดคือ เทคนิคเพื่อนบ้านใกล้ที่สุด ซึ่งมีค่าเท่ากับ 76.84% รองลงมาคือ เทคนิคนาอิวเบย์ มีเท่ากับ 48.54% และเทคนิคที่ให้ค่าเฉลี่ยความไววน้อยที่สุดคือ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน ซึ่งมีค่าเท่ากับ 26.30% และเทคนิคที่ให้ค่าเฉลี่ยความจำเพาะสูงสุดคือ เทคนิคต้นไม้ป่าสุ่ม ซึ่งมีค่าเท่ากับ 99.88% รองลงมาคือ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน มีเท่ากับ 98.74% และเทคนิคที่ให้ค่าเฉลี่ยความจำเพาะน้อยที่สุดคือ เทคนิคนาอิวเบย์ ซึ่งมีค่าเท่ากับ 83.30% และผู้วิจัยได้นำค่าเฉลี่ยของค่าความแม่นยำของแต่ละเทคนิคมาสร้างเป็นแผนภูมิแท่งเพื่อเปรียบเทียบประสิทธิภาพของตัวแบบ ดังแสดงใน Figure 5

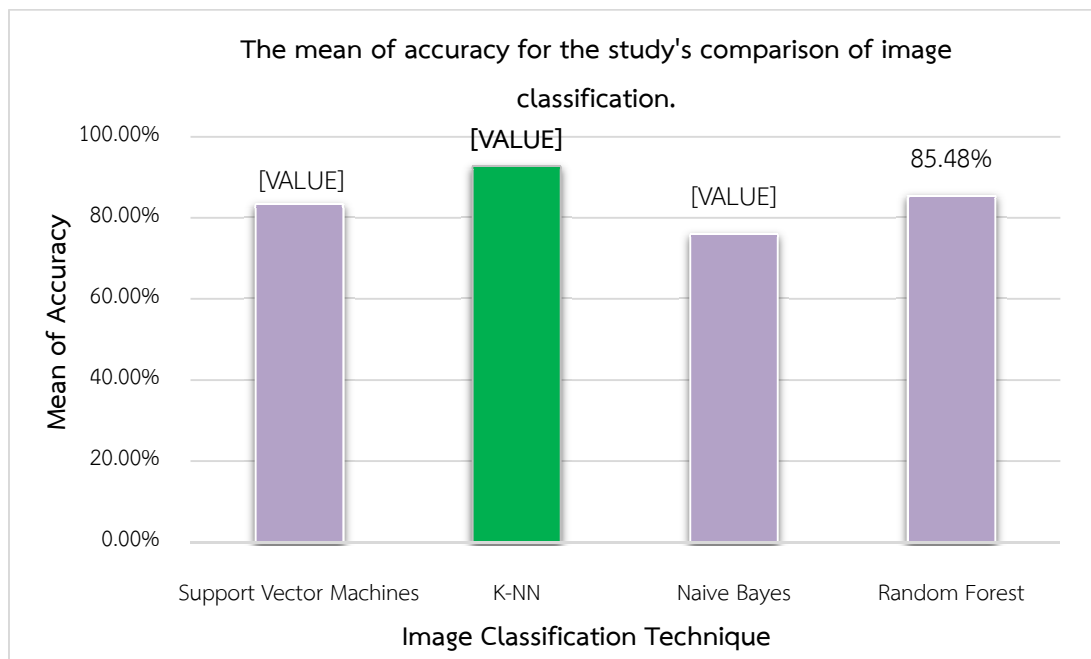


Figure 5 Comparing the Performance of Image Classification Models

จาก Figure 5 แสดงการเปรียบเทียบประสิทธิภาพของตัวแบบการจำแนกประเภทข้อมูลภาพ พบว่า เทคนิคเพื่อนบ้านใกล้ที่สุด มีค่าประสิทธิภาพของการจำแนกประเภทข้อมูลภาพที่ดีที่สุด โดยมีค่าเฉลี่ยของความแม่นยำสูงสุด ซึ่งมีค่าเท่ากับ 92.64% รองลงมาคือ เทคนิคต้นไม้ป่าสุ่ม มีค่าเฉลี่ยความแม่นยำเท่ากับ 85.48% และเทคนิคที่ให้ค่าเฉลี่ยความแม่นยำน้อยที่สุดคือ เทคนิคนาอิวเบย์ ซึ่งมีค่าเท่ากับ 75.88% ดังนั้น เทคนิคเพื่อนบ้านใกล้ที่สุดจึงมีความเหมาะสมที่สุดสำหรับนำไปพัฒนาเพื่อสร้างตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในต่อไป

4. สรุปผลการวิจัย

จากการทดลองสร้างตัวแบบและเปรียบเทียบประสิทธิภาพของตัวแบบสำหรับการคัดกรองผู้ป่วยโรคนี้ในไตด้วยภาพถ่าย CT scan โดยใช้ข้อมูลจริงของผู้ป่วยที่ได้เก็บรวบรวมไว้ในปี พ.ศ.2565 จำนวนทั้งสิ้น 6,454 ภาพ ซึ่งประกอบด้วยข้อมูลภาพถ่าย CT scan ของผู้ป่วยโรคนี้ในไต (S) และข้อมูลภาพถ่าย CT scan ของคนปกติ (N) และได้วิเคราะห์ตามกระบวนการมาตรฐานการทำเหมืองข้อมูล โดยใช้เทคนิคการจำแนกประเภทข้อมูลภาพ จำนวน 4 เทคนิค ได้แก่ เทคนิคซัพพอร์ตเวกเตอร์แมชชีน เทคนิคเพื่อนบ้านใกล้ที่สุด เทคนิคนาอิวเบย์ และเทคนิคต้นไม้ป่าสุ่ม จากนั้นทำการเปรียบเทียบประสิทธิภาพของการจำแนกประเภทข้อมูลด้วยค่าความแม่นยำ ค่าประสิทธิภาพโดยรวม ค่าความไว และค่าจำเพาะ โดยใช้โปรแกรม RapidMiner Studio Version 10 สำหรับการเตรียมข้อมูล สร้างตัวแบบ และการวิเคราะห์ข้อมูล เมื่อทำการวิเคราะห์ข้อมูล พบว่า เทคนิคเพื่อนบ้านใกล้ที่สุด เป็นเทคนิคที่เหมาะสมสำหรับนำมาสร้างตัวแบบการคัดกรองผู้ป่วยโรคนี้ในไตโดยใช้ภาพถ่าย CT scan เนื่องจากให้ค่าประสิทธิภาพของการจำแนกประเภทข้อมูลภาพที่สูงที่สุด โดยมีค่าเฉลี่ยของค่าความแม่นยำเท่ากับ 92.64% ค่าประสิทธิภาพโดยรวมเท่ากับ 80.72% ค่าความไวเท่ากับ 76.84% และค่าจำเพาะเท่ากับ 96.34% ซึ่งสอดคล้องงานวิจัยของ Verma, Nath, Tripathi, and Saini (2017) ที่ได้ศึกษาเรื่องการวิเคราะห์และการระบุในไต โดยใช้เทคนิคเพื่อนบ้านใกล้ที่สุดและเทคนิคซัพพอร์ตเวกเตอร์แมชชีน ที่มีค่าความแม่นยำของอัลกอริทึมมากกว่า 85% ซึ่งผลลัพธ์ที่ได้จากการวิเคราะห์ข้อมูลในครั้งนี้ยังสามารถนำตัวแบบที่ได้ไปพัฒนาเป็นระบบสารสนเทศสำหรับการคัดกรองความเสี่ยงผู้ป่วยโรคนี้ในไตด้วยภาพเบื้องต้น เพื่อช่วยลดภาระของทีมแพทย์และสหวิชาชีพในการคัดกรองผู้ป่วยโรคนี้ในไตได้อีกทางหนึ่ง และทำให้การวินิจฉัยมีความรวดเร็วขึ้น

ผลประโยชน์ทับซ้อน

ผู้เขียนขอยืนยันว่างานวิจัยนี้ไม่มีความขัดแย้งทางผลประโยชน์ที่เกี่ยวข้องกับบทความนี้กับบทความอื่น ๆ

REFERENCES

- Laongkum, A., and Kumeechai, P. (2020). Implementation of classification system for Buddha amulet using GLCM and Wavelet Transform and using Neural Network for classify. *Royal Thai Naval Academy Journal of Science and Technology*, 3(1), 9-20. [In Thai].
- Sukprasert, A. (2021). *Data Mining with RapidMiner Studio*. Mahasarakham University. Mahasarakham Business School. [In Thai].
- Promkasikorn, L. (2015). *The factors predicting the health behaviors related to the occurrence of kidney stone formation among kidney stone disease patients*. (Master's thesis). Thammasat University, Faculty of Nursing, Department of Adult Nursing and the Aged. [In Thai].
- MD. Mehedi, K, H. (2022). CT kidney dataset: normal-cyst-tumor and stone. Retrieved from <https://www.kaggle.com/datasets/nazmul0087/ct-kidney-dataset-normal-cyst-tumor-and-stone>.
- Nonsiri, N., Chaichitwanidchakol, P., and Somkantha, K. (2022). Data classification for diabetes risk diagnosis using majority voting ensemble method and forward feature selection method. *Udon Thani Rajabhat University Journal of Sciences and Technology*, 10(2), 107–122. [In Thai].
- Qadri, S. (2021). Role of machine vision for identification of kidney stones using multi features analysis. *Lahore Garrison University Research Journal of Computer Science and Information Technology*, 5(3), 1-14.
- 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, 574-580.
- Pakkasung, Y., Methakanjanasak, N., and Theeranut, A. (2019). Illness perception of patient with recurrent renal stone and chronic kidney disease. *Journal of Nursing and Health Care*, 37(1), 249-257. [In Thai].