

การวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลว
ด้วยวิธีการค้นหาเพื่อนบ้านใกล้ที่สุด

DIAGNOSING THE RISK OF DEATH FROM HEART FAILURE
BY K-NEAREST NEIGHBOR (K-NN)

มุกิตา บุตรลม อริษา ขุนนนท์ และ วิไลพร กุลตังวัฒนา *

สาขาวิชาวิทยาการคอมพิวเตอร์และเทคโนโลยีสารสนเทศ คณะวิทยาศาสตร์ มหาวิทยาลัยราชภัฏอุดรธานี

Mutita Butlom Alisa Kunnon and Wilaiporn Kultangwattana *

Department of Computer Science and Information Technology, Faculty of Science,
Udon thani Rajchabhat University

(Received: January 9, 2023; Revised: April 15, 2023; Accepted: April 20, 2023)

*ผู้ประสานงาน : ผู้ช่วยศาสตราจารย์วิไลพร กุลตังวัฒนา อีเมล : wilaiporn.ku@udru.ac.th

บทคัดย่อ

งานวิจัยนี้เสนอระบบการวินิจฉัยภาวะความเสี่ยงในการ เสียชีวิตจากโรคหัวใจล้มเหลวด้วยวิธีการค้นหาเพื่อนบ้านใกล้ที่สุด โดย พิจารณาจากคุณลักษณะจาก อายุ,โรคหิตจาง,ระดับของเอนไซม์ในเลือด, โรคเบาหวาน,ร้อยละของเลือดที่ออกจากหัวใจ,ความดันโลหิตสูง,เกล็ด เลือด,ระดับครีเอตินินในเลือด,ระดับโซเดียมในเลือด,เพศ,การสูบบุหรี่,เวลา ที่ซึ่งมีจำนวน 12 คุณลักษณะเด่น โดยคุณลักษณะเด่นทั้งหมดจะเป็นข้อมูลเข้าสำหรับวิธีการเพื่อนบ้านที่ใกล้เคียงที่สุดจำนวน K ตัว (K-Nearest Neighbor) เพื่อการฝึกและการทดสอบ โดยจำแนกข้อมูลออกเป็น เสี่ยง และไม่เสี่ยง การทดสอบประสิทธิภาพของระบบการวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลวจากวิธีการเพื่อนบ้านใกล้ที่สุด ทดสอบด้วยวิธี 10-fold Cross validation จากผลการทดลองพบว่าระบบสามารถวิเคราะห์การวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลวมีความถูกต้องอยู่ที่ 83.17% จากผลงานวิจัยนี้เชื่อว่าวิธีการที่นำเสนอสามารถเป็นเครื่องมือในการวินิจฉัยความเสี่ยงในการเสียชีวิตจากโรคหัวใจอย่างมีประสิทธิภาพ

คำสำคัญ: อัลกอริทึมเพื่อนบ้านใกล้ที่สุด, โรคหัวใจล้มเหลว, การจำแนก, เหมืองข้อมูล

ABSTRACT

This research proposes a system for diagnosing the risk of death from heart failure by the nearest neighbor search method. The proposed method used characteristics from Age, anemia, blood enzyme levels, diabetes, cardiac output percentage, high blood pressure, platelets, serum creatinine, serum sodium, sex, smoke, and time which have 12 features. All features are taken as input into K-Nearest Neighbor algorithm for training and testing for diagnosing the Risk of Death from Heart Failure by K-Nearest Neighbor (K-NN). The results were classified as risky and non-risk. In this paper, the 10-fold cross validation is used for testing the efficiency of the proposed method. From the experimental results show that the proposed method had accurate about 83.14%. We can believe that the proposed method can be a tool for diagnosing the risk of death from heart failure efficiently.

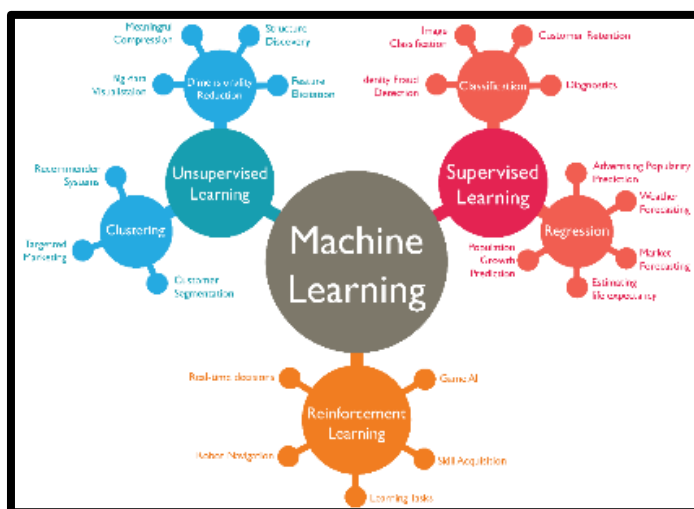
Keywords: Nearest Neighbor Algorithm, Heart Failure, classification

1. บทนำ

ภาวะหัวใจล้มเหลวหรือหัวใจวาย (Heart Failure) เป็นภาวะที่เกิดจากความผิดปกติของการทำงานของหัวใจ ซึ่งเป็นผลมาจากกล้ามเนื้ออ่อนแรง หรือขาดความยืดหยุ่นส่งผลให้ไม่สามารถสูบฉีดเลือดไปเลี้ยงที่อวัยวะต่าง ๆ ของร่างกายได้เพียงพอจึงทำให้อวัยวะขาดออกซิเจนไปหล่อเลี้ยงและหัวใจจะไม่สามารถรับเลือดกลับเข้าสู่หัวใจได้ สาเหตุที่พบได้บ่อยเกิดจากโรคกล้ามเนื้อหัวใจขาดเลือดจากหลอดเลือดหัวใจตีบ หรือเกิดหัวใจวาย [1] สาเหตุที่พบได้บ่อยเกิดจากโรคกล้ามเนื้อหัวใจขาดเลือดจากหลอดเลือดหัวใจตีบ หรือเกิดหัวใจวาย และภาวะหัวใจล้มเหลว สาเหตุอื่นที่พบ ได้แก่ ความดันโลหิตสูง โรคลิ้นหัวใจ และการติดเชื้อของกล้ามเนื้อหัวใจ สาเหตุต่าง ๆ เหล่านี้ ทำให้กล้ามเนื้อหัวใจสูบฉีดเลือดลดลง ทำให้ปริมาณเลือดไหลผ่านร่างกายลดลงด้วย หัวใจจึงพยายามปรับตัวให้ปริมาณเลือดซึ่งจะสูบฉีดไปเลี้ยงส่วนต่าง ๆ ของร่างกายเป็นปกติ ห้องหัวใจจึงขยายตัวเพื่อเก็บกักเลือดให้มากขึ้น ซึ่งจะทำให้ได้เพียงระยะหนึ่งเท่านั้น เพราะกล้ามเนื้อห้องหัวใจที่ยืดขยายนี้จะอ่อนล้าและไม่สามารถสูบฉีดได้ดีอีกต่อไป ในการบีบตัวแต่ละครั้ง หัวใจที่อ่อนกำลังจะสูบฉีดเลือดได้น้อยลง [2] ดังนั้นเราจึงพัฒนาระบบนี้เพื่อช่วยประเมินความเสี่ยงของผู้ที่เป็นภาวะโรคหัวใจล้มเหลวมีความเสี่ยงที่จะเสียชีวิตหรือไม่ เพื่อที่จะได้ดูแลรักษาอย่างถูกวิธีจากแพทย์ผู้เชี่ยวชาญ

2. ทฤษฎีและงานวิจัยที่เกี่ยวข้อง

การเรียนรู้ของเครื่องจักร (machine learning) เป็นการศึกษาอัลกอริทึมของคอมพิวเตอร์ที่มีการพัฒนาการเรียนรู้ของเครื่องถูกมองว่าเป็นส่วนหนึ่งของปัญญาประดิษฐ์ โดยอัลกอริทึมสร้างแบบจำลองทางคณิตศาสตร์จากข้อมูลตัวอย่าง (เรียกว่า ข้อมูลสอน) เพื่อที่จะคาดการณ์หรือตัดสินใจได้อย่างชัดเจน การเรียนรู้ของเครื่องพัฒนามาจากการศึกษาการรู้จำแบบ เกี่ยวข้องกับการศึกษาและการสร้างอัลกอริทึมที่สามารถเรียนรู้ข้อมูลและทำนายข้อมูลได้ [9]. อัลกอริทึมนั้นจะทำงานโดยอาศัยโมเดลที่สร้างมาจากชุดข้อมูลตัวอย่างเข้าเพื่อการทำนายหรือตัดสินใจในภายหลัง แทนที่จะทำงานตามลำดับของคำสั่งโปรแกรมคอมพิวเตอร์ การเรียนรู้ของเครื่องมีเกี่ยวข้องอย่างมากกับสถิติศาสตร์ เนื่องจากทั้งสองสาขาศึกษาการวิเคราะห์ข้อมูลเพื่อการทำนายเช่นกัน นอกจากนี้ยังมีความสัมพันธ์กับสาขาการหาค่าเหมาะที่สุดในทางคณิตศาสตร์ที่แบ่งของวิธีการ ทฤษฎี และการประยุกต์ใช้ การเรียนรู้ของเครื่องสามารถนำไปประยุกต์ใช้งานได้หลากหลาย ไม่ว่าจะเป็นการกรองอีเมลขยะ การรู้จำตัวอักษร เครื่องมือค้นหา และคอมพิวเตอร์วิทัศน์ [7]. Machine learning สามารถจำแนกได้ 3 ประเภทใหญ่ๆ คือ การเรียนรู้แบบมีผู้สอน (supervised learning) การเรียนรู้แบบไม่มีผู้สอน (unsupervised learning) และ การเรียนรู้แบบเสริมแรง (Reinforcement Learning) ซึ่งงานของเราได้นำการเรียนรู้แบบผู้สอน (Supervised Learning) มาใช้ K-Nearest Neighbor [3]



รูปที่ 1 อัลกอริทึมของ Machine Learning

K-Nearest Neighbor (K-NN) จัดว่าเป็นการจำแนกแบบมีผู้ฝึกสอน (Supervised Machine Learning Algorithm) [5]. K-NN เป็นวิธีการแบ่งคลาสสำหรับใช้จัดหมวดหมู่ข้อมูล (Classification) [8] ใช้หลักการเปรียบเทียบข้อมูลที่สนใจกับข้อมูลอื่นว่ามีความคล้ายคลึงมากน้อยเพียงใด หากข้อมูล

ที่กำลังสนใจนั้นอยู่ใกล้ข้อมูลใดมากที่สุด ระบบจะให้ คำตอบเป็นเหมือนคำตอบของข้อมูลที่อยู่ใกล้ที่สุดนั้นลักษณะการทำงานแบบไม่ได้ใช้ข้อมูลชุดเรียนรู้ (training data) ในการสร้างแบบจำลองแต่จะใช้ข้อมูลนั้นมาเป็นตัวแบบจำลอง [6]

$$d(p, q) = \sqrt{\sum_{i=1}^n (q_i - p_i)^2} \quad (1)$$

d คือระยะทางจากจุด p ไปยังจุด q

p ข้อมูลกลุ่มตัวอย่าง หรือ training data ที่นำมาพิจารณา

q ข้อมูลที่ต้องการจำแนก

n คือจำนวนมิติของข้อมูล

ตัวอย่างขั้นตอนการจำแนกข้อมูล K-Nearest Neighbor

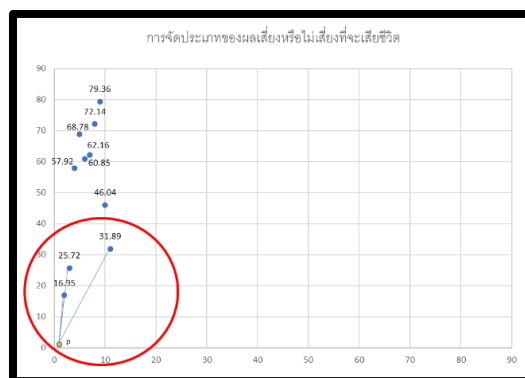
ตารางที่ 1 ชุดข้อมูลภาวะโรคหัวใจล้มเหลว

No	อายุ	โรค หัวใจ	ระดับของเอนไซม์ ในเลือด	โรคเบา หวาน	ร้อยละของเลือด ที่ออกจากหัวใจ	ความดัน โลหิตสูง	เกล็ด เลือด	ระดับครีเอตินีน ในเลือด	ระดับ โซเดียม ในเลือด	เพศ	การสูบบุหรี่	เวลา ติดตาม	ความ เสี่ยง
1	42	0	0.07	0	60	0	0.48	1.18	0.7	0	0	82	0
2	45	0	0.07	0	80	0	0.48	1.18	0.7	0	0	63	0
3	70	0	0.07	0	20	1	0.48	1.83	0.6	1	1	31	1
4	94	0	0.07	1	38	1	0.48	1.83	0.6	1	0	27	1
5	75	0	0.07	1	30	1	0.48	1.83	0.6	0	0	23	1
6	65	0	0.02	0	65	0	0.48	1.5	0.7	0	0	10	1
7	87	1	0.02	0	38	0	0.48	0.9	0.8	1	0	14	1
8	75	0	0.07	0	20	1	0.48	1.9	0.5	1	0	4	1
9	69	0	0.07	0	20	0	0.48	1.2	0.6	1	1	73	1
10	60	0	0.5	1	62	0	0.49	6.8	1	0	0	43	1

New instance (อายุ = 40, โรคหัวใจ = 1, ระดับของเอนไซม์ในเลือด = 0.09, โรคเบาหวาน = 0, ร้อยละของเลือดที่ออกจากหัวใจ = 55, ความดันโลหิตสูง = 1, เกล็ดเลือด = 0.45, ระดับครีเอตินีนในเลือด = 0.8, ระดับโซเดียมในเลือด = 1.18, เพศ = 0, การสูบบุหรี่ = 0, เวลาติดตาม = 66) คำนวณระยะห่างระหว่างข้อมูลที่ต้องการจำแนกกับชุดข้อมูลสอน (training data) ทั้งหมด เพื่อวัดความคล้ายคลึงของข้อมูลด้วย Euclidean Distance [10]

$$\begin{aligned}
 d(q, p1) &= \sqrt{(40-42)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-0)^2 + (55-60)^2 + (1-0)^2 + (0.45-0.48)^2 + (0.8-1.18)^2 + (1.18-0.7)^2 + (0-0)^2 + (0-0)^2 + (66-82)^2} = 16.95 \\
 d(q, p2) &= \sqrt{(40-45)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-0)^2 + (55-80)^2 + (1-0)^2 + (0.45-0.48)^2 + (0.8-1.18)^2 + (1.18-0.7)^2 + (0-0)^2 + (0-0)^2 + (66-63)^2} = 25.72 \\
 d(q, p3) &= \sqrt{(40-70)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-0)^2 + (55-20)^2 + (1-1)^2 + (0.45-0.48)^2 + (0.8-1.83)^2 + (1.18-0.6)^2 + (0-1)^2 + (0-1)^2 + (66-31)^2} = 57.92 \\
 d(q, p4) &= \sqrt{(40-94)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-1)^2 + (55-38)^2 + (1-1)^2 + (0.45-0.48)^2 + (0.8-1.83)^2 + (1.18-0.6)^2 + (0-1)^2 + (0-0)^2 + (66-27)^2} = 68.78 \\
 d(q, p5) &= \sqrt{(40-75)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-1)^2 + (55-30)^2 + (1-1)^2 + (0.45-0.48)^2 + (0.8-1.83)^2 + (1.18-0.6)^2 + (0-0)^2 + (0-0)^2 + (66-23)^2} = 60.85 \\
 d(q, p6) &= \sqrt{(40-65)^2 + (1-0)^2 + (0.09-0.02)^2 + (0-0)^2 + (55-65)^2 + (1-0)^2 + (0.45-0.48)^2 + (0.8-1.5)^2 + (1.18-0.7)^2 + (0-0)^2 + (0-0)^2 + (66-10)^2} = 62.16 \\
 d(q, p7) &= \sqrt{(40-75)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-1)^2 + (55-38)^2 + (1-0)^2 + (0.45-0.48)^2 + (0.8-0.9 + (1.18-0.8)^2 + (0-1)^2 + (0-0)^2 + (66-14)^2} = 72.14 \\
 d(q, p8) &= \sqrt{(40-75)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-0)^2 + (55-20)^2 + (1-1)^2 + (0.45-0.48)^2 + (0.8-1.59 + (1.18-0.5)^2 + (0-1)^2 + (0-0)^2 + (66-4)^2} = 79.36 \\
 d(q, p9) &= \sqrt{(40-69)^2 + (1-0)^2 + (0.09-0.07)^2 + (0-0)^2 + (55-20)^2 + (1-0 + (0.45-0.48)^2 + (0.8-1.2)^2 + (1.18-0.6)^2 + (0-1)^2 + (0-1)^2 + (66-73)^2} = 46.04 \\
 d(q, p10) &= \sqrt{(40-60)^2 + (1-0)^2 + (0.09-0.5)^2 + (0-1)^2 + (55-62)^2 + (1-0)^2 + (0.45-0.49)^2 + (0.8-6.8)^2 + (1.18-1)^2 + (0-0)^2 + (0-0)^2 + (66-43)^2} = 31.89
 \end{aligned}$$

นับจำนวนตัวอย่างที่ใกล้เคียงกับข้อมูลที่ต้องการทราบที่สุด K ตัว และทำการจำแนกประเภท กำหนด $k=3$



รูปที่ 2 การจัดประเภทของผลเสียหรือไม่เสียที่จะเสียชีวิต

กำหนด $k=3$ ดังนั้นค่าที่ใกล้เคียงที่สุด 3 ค่าคือ

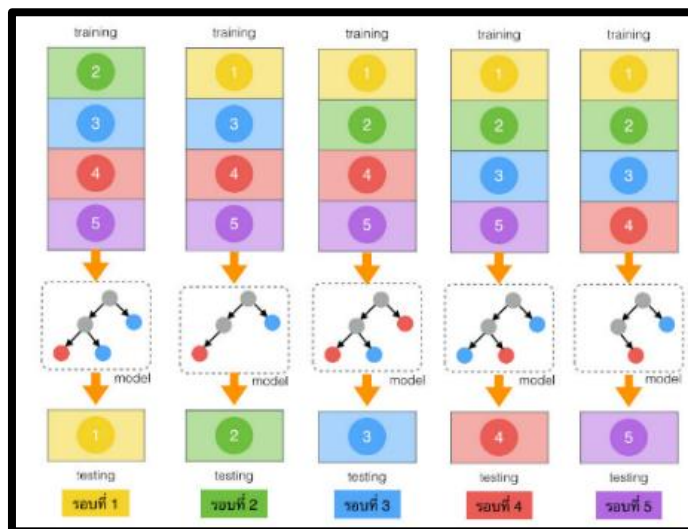
$$d(q, p1) = 16.95$$

$$d(q, p2) = 25.72$$

$$d(q, p10) = 31.89$$

การทำนายข้อมูล ชุดใหม่ถูกจำแนกให้อยู่ในกลุ่มไม่เสี่ยง 2 ตัวอย่าง และอยู่ในกลุ่มเสี่ยงเพียง 1 ตัวอย่าง ดังนั้นผลการจำแนกคือ จัดอยู่ในประเภทไม่เสี่ยงที่จะเสียชีวิต

Cross-validation Test วิธีนี้เป็นวิธีที่นิยมในการทำงานวิจัย เพื่อใช้ในการทดสอบประสิทธิภาพของโมเดลเนื่องจากผลที่ได้มีความน่าเชื่อถือ การวัดประสิทธิภาพด้วยวิธี Cross-validation จะทำการแบ่งข้อมูลออกเป็นหลายส่วน เช่น 5-fold cross-validation คือ ทำการแบ่งข้อมูลออกเป็น 5 ส่วน โดยที่แต่ละส่วนมีจำนวนข้อมูลเท่ากัน หรือ 10-fold cross-validation คือ การแบ่งข้อมูลออกเป็น 10 ส่วน โดยที่แต่ละส่วนมีจำนวนข้อมูลเท่ากัน หลังจากนั้นข้อมูลหนึ่งส่วนจะใช้เป็นตัวทดสอบประสิทธิภาพของโมเดล ทำวนไปเช่นนี้จนครบจำนวนที่แบ่งไว้ [4]



รูปที่ 3 วิธี Cross-validation Test

3. วิธีดำเนินการวิจัย

การดำเนินงานวิจัยระบบการวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลว ด้วยวิธีการค้นหาเพื่อนบ้านใกล้เคียงที่สุดแบ่งออกเป็น 2 ขั้นตอน คือ 1) ขั้นตอนการเตรียมข้อมูล 2) ขั้นตอนการทดสอบประสิทธิภาพ

3.1 การเตรียมข้อมูล

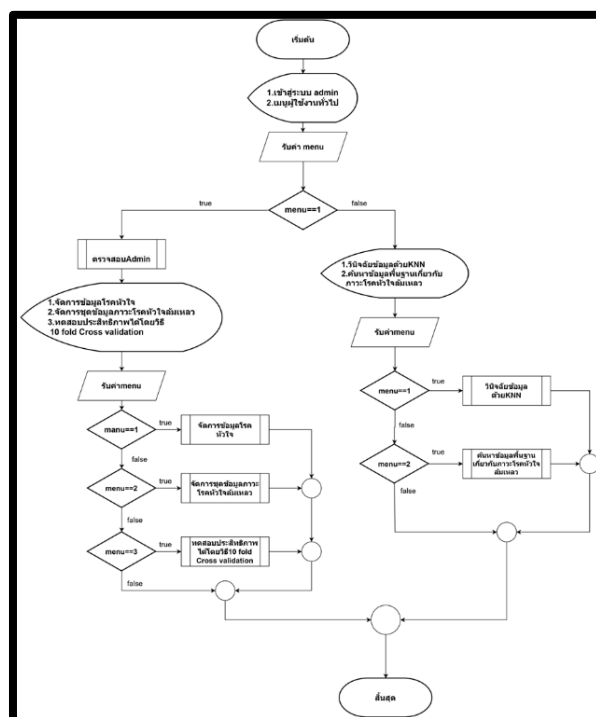
ข้อมูลที่ใช้เป็นข้อมูลที่ได้จากเว็บไซต์ Kaggle ทั้งหมด 100 ตัวอย่าง เป็นเสี่ยง 67 และไม่เสี่ยง 33 โดยมีลักษณะทั้งหมด 12 คุณลักษณะ มีดังนี้ อายุ, โรคหัวใจ, ระดับของเอนไซม์ในเลือด, โรคเบาหวาน, ร้อยละของเลือดที่ออกจากหัวใจ, ความดันโลหิตสูง, เกล็ดเลือด, ระดับครีเอตินินในเลือด, ระดับโซเดียมในเลือด, เพศ, การสูบบุหรี่, เวลา

anaemia	creatinine_phosphokinase	diabetes	ejection_fraction	high_blood_pressure	platelets	serum_creatinine	serum_sodium	sex	smoking	time	DEATH_EVENT
โรคไตเรื้อรัง	ระดับของเอนไซม์ในเลือด	โรคเบาหวาน	อัตราการบีบตัวของหัวใจ	ความดันโลหิตสูง	เกล็ดเลือด	ระดับเอนไซม์ในเลือด	ระดับโซเดียมในเลือด	เพศ	สูบบุหรี่	ระยะเวลาติดตามผล	
0	0.07	0	60	0	0.48	1.18	0.7	0	0	82	0
0	0.07	0	80	0	0.48	1.18	0.7	0	0	63	0
0	0.07	0	20	1	0.48	1.83	0.6	1	1	31	1
0	0.07	1	38	1	0.48	1.83	0.6	1	0	27	1
0	0.07	1	30	1	0.48	1.83	0.6	0	0	23	1
0	0.02	0	65	0	0.48	1.5	0.7	0	0	10	1
1	0.02	0	38	0	0.48	0.9	0.8	1	0	14	1
0	0.07	0	20	1	0.48	1.9	0.5	1	0	4	1
0	0.07	0	20	0	0.48	1.2	0.6	1	1	73	1
0	0.50	1	62	0	0.49	6.8	1.0	0	0	43	1
0	1.00	0	38	0	0.48	1.1	0.7	1	0	6	1
0	0.17	0	25	1	0.48	0.9	0.5	1	0	38	1
1	0.03	1	38	0	0.50	2.2	0.5	0	1	45	1
1	0.00	0	60	0	0.46	1.1	0.7	0	0	85	0
1	0.06	1	25	1	0.46	1.3	0.6	1	0	83	0
0	0.04	1	20	1	0.46	1.3	0.7	1	1	59	1
0	0.01	0	35	0	0.46	1.1	0.8	1	1	60	0
0	0.03	0	25	1	0.46	0.9	0.8	1	1	10	1
1	0.01	0	25	0	0.46	1	0.8	1	1	65	1
1	0.02	0	38	1	0.50	1.1	0.7	1	0	11	1
1	0.00	0	25	0	0.51	1.3	0.7	0	0	16	0
0	0.01	1	45	0	0.51	0.8	0.7	1	1	80	0
0	0.02	0	25	0	0.45	1.2	0.9	0	0	66	1
1	0.02	1	38	1	0.44	0.6	0.5	1	1	74	0
0	0.01	1	45	1	0.52	1.3	0.7	1	1	26	1
0	0.75	0	35	0	0.53	1	0.8	1	1	72	1
1	0.03	0	35	1	0.44	0.9	0.8	1	1	20	1
1	0.01	0	25	1	0.54	1	0.8	0	0	15	1

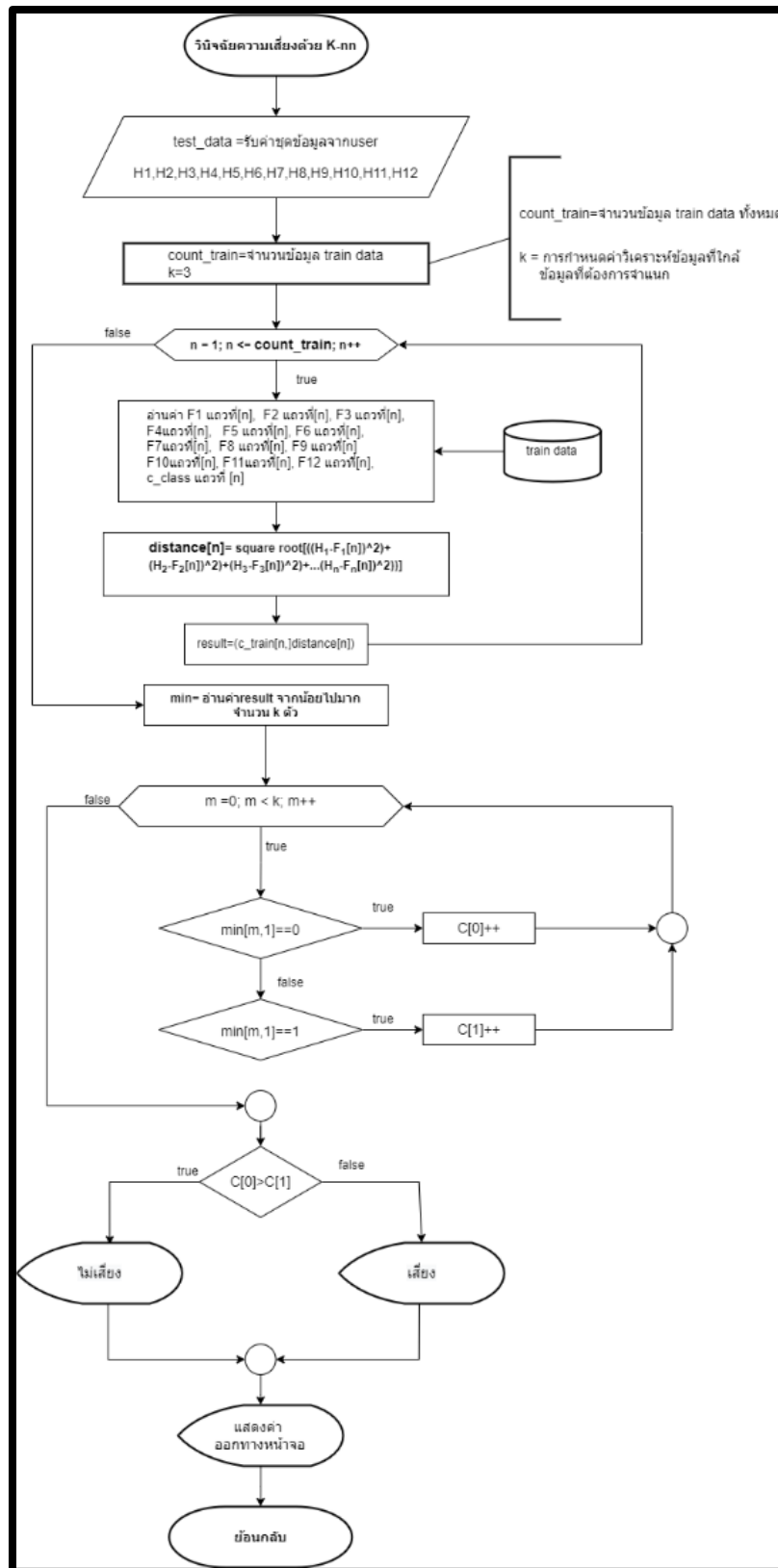
รูปที่ 4 ภาพตัวอย่างข้อมูลที่นำมาใช้ในการวินิจฉัย

3.2 การทดสอบประสิทธิภาพ

ในงานวิจัยนี้ผู้วิจัยได้ทำการทดสอบประสิทธิภาพของระบบการวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลวด้วยวิธีการค้นหาเพื่อนบ้านใกล้เคียงด้วยวิธี 10-fold Cross validation พบว่าระบบสามารถวิเคราะห์การวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตโรคหัวใจล้มเหลว มีความถูกต้องอยู่ที่ 83.17%



รูปที่ 5 ผังระบบงาน System Flow Chart



รูปที่ 6 Flow Chart วินิจฉัยข้อมูลด้วย KNN

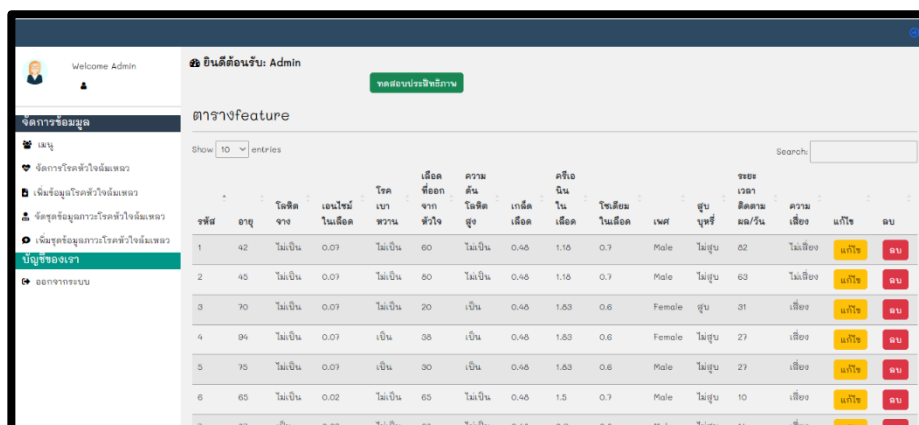
4. ผลการศึกษา

จากการศึกษาและพัฒนากระบวนการวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลวด้วยวิธีการค้นหาเพื่อนบ้านใกล้ที่สุด โดยระบบนี้มีผู้ใช้งาน 2 แบบ คือ ผู้ดูแลระบบ ผู้ใช้งานทั่วไปและผลทดสอบประสิทธิภาพมีความถูกต้อง 83.17%

4.1 ส่วนของผู้ดูแลระบบ



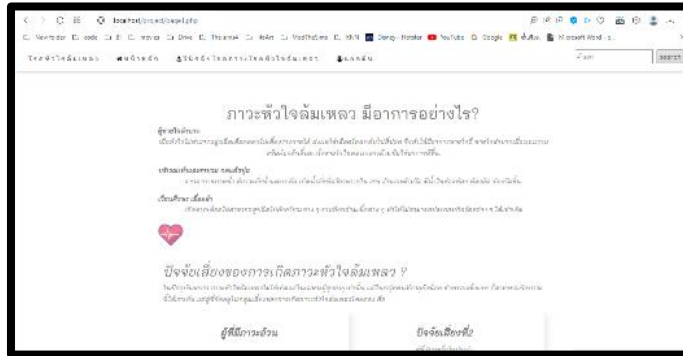
รูปที่ 7 หน้าจอจัดการข้อมูลโรคหัวใจล้มเหลว (ผู้ดูแลระบบ)



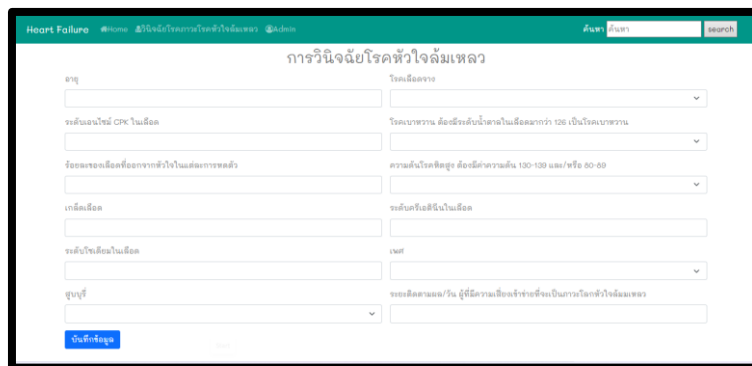
ID	ชื่อ	อายุ	เพศ	ความดันโลหิต	คอเลสเตอรอล	น้ำตาลในเลือด	อัตราการเต้นของหัวใจ	ระยะเวลาในการรักษา	ความรุนแรง	สถานะ
1	ชาย	42	ชาย	120/80	180	100	70	10	เบา	ลบ
2	ชาย	45	ชาย	130/85	190	110	75	12	เบา	ลบ
3	ชาย	70	ชาย	140/90	200	120	80	15	เบา	ลบ
4	ชาย	94	ชาย	150/95	210	130	85	18	เบา	ลบ
5	ชาย	75	ชาย	160/100	220	140	90	20	เบา	ลบ
6	ชาย	65	ชาย	170/105	230	150	95	22	เบา	ลบ

รูปที่ 8 หน้าจอจัดการชุดข้อมูลภาวะโรคหัวใจล้มเหลว (ผู้ดูแลระบบ)

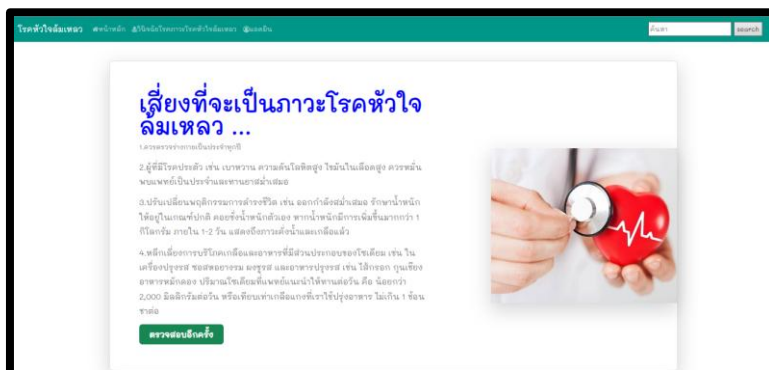
4.2 ส่วนของผู้ใช้งานทั่วไป



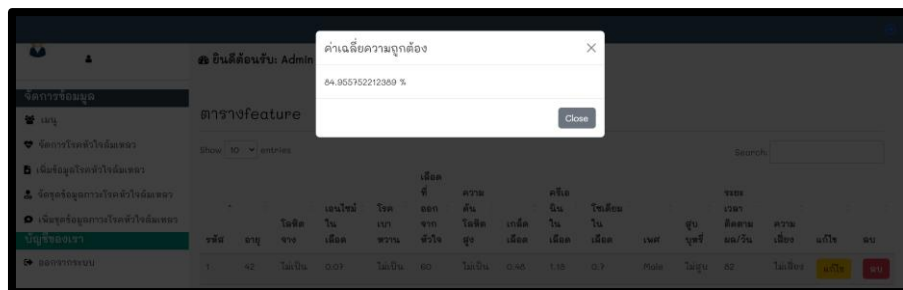
รูปที่ 9 หน้าจอความรู้พื้นฐานเกี่ยวกับโรคหัวใจล้มเหลว (ผู้ใช้งาน)



รูปที่ 10 หน้าจอการวินิจฉัยโรคหัวใจล้มเหลว (ผู้ใช้งาน)



รูปที่ 11 หน้าจอแสดงการวินิจฉัยโรคหัวใจล้มเหลว (ผู้ใช้งาน)



รูปที่ 12 หน้าจอแสดงผลการทดสอบประสิทธิภาพ

ตารางที่ 2 สรุปผลการประเมินประสิทธิภาพของระบบจากผู้ใช้ระบบงาน

รายการประเมิน	$\bar{X}$	S. D	ผลแปล
1. การออกแบบมีความเหมาะสมกับการเป็นเว็บไซต์การวินิจฉัย	4.2	0.63	ดี
2. การแสดงข้อมูลและการวินิจฉัย	4.4	0.52	ดี
3. ความสวยงามของเว็บไซต์	4.5	0.53	ดีมาก
4. ความยากง่ายในการใช้งานของเว็บไซต์	4.4	0.52	ดี
5. ความสามารถของเว็บไซต์ตรงกับความต้องการของผู้ใช้งาน	4.4	0.52	ดี
6. การประมวลผลของระบบ	4.2	0.63	ดี
ผลรวม	4.35	0.64	ดี

จากตารางที่ 2 การประเมินประสิทธิภาพของระบบในส่วนผู้ใช้งานระบบจากคนประเมิน 10 คนอยู่ในระดับ ดี ($\bar{X}$ = 4.35, S.D = 0.64)

ตารางที่ 3 สรุปผลการประเมินประสิทธิภาพของระบบจากผู้ดูแลระบบงาน

รายการประเมิน	$\bar{X}$	S. D	ผลแปล
1. ความยากง่ายในการใช้งานของระบบ	4.8	0.44	ดีมาก
2. การประมวลผลของระบบ	4.8	0.44	ดีมาก
3. ความสวยงามของระบบ	4.2	0.44	ดี
4. ความปลอดภัยของระบบ	4.6	0.54	ดีมาก
ผลรวม	4.6	0.50	ดีมาก

จากตารางที่ 3 การประเมินประสิทธิภาพของระบบในส่วนผู้ดูแลระบบจากผลประเมิน 5 คนอยู่ในระดับ ดีมาก ($\bar{X}$ = 4.6, S.D = 0.50)

5.บทสรุปและข้อเสนอแนะ

ระบบการวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลวด้วยวิธีการค้นหาเพื่อนบ้านใกล้เคียงที่สุด ได้พัฒนาขึ้นเพื่อช่วยประเมินความเสี่ยงของผู้ที่เป็นภาวะโรคหัวใจล้มเหลวมีความเสี่ยงที่จะเสียชีวิตหรือไม่ เพื่อที่จะได้ดูแลรักษาอย่างถูกวิธีจากแพทย์ผู้เชี่ยวชาญ โดยระบบนี้มีผู้ใช้งาน 2 แบบ คือผู้ดูแลระบบและผู้ใช้ทั่วไป ซึ่งผู้ดูแลระบบสามารถจัดการเพิ่ม ลบ แก้ไข ข้อมูลพื้นฐานโรคหัวใจและชุดข้อมูลภาวะโรคหัวใจล้มเหลวได้ ส่วนผู้ใช้งานทั่วไปสามารถดูข้อมูลพื้นฐานเกี่ยวกับภาวะโรคหัวใจล้มเหลวได้ ทำแบบคัดกรองว่ามีความเสี่ยงที่จะเสียชีวิตจากโรคหัวใจล้มเหลว ดูผลประเมินแบบคัดกรอง และยังสามารถดูผลทดสอบประสิทธิภาพได้ ผลทดสอบประสิทธิภาพมีความถูกต้อง 83.17% หากมีการพัฒนาควรพัฒนาระบบการวินิจฉัยภาวะความเสี่ยงในการเสียชีวิตจากโรคหัวใจล้มเหลวเป็น Mobile Application เพื่อให้สามารถใช้งานให้สะดวกและง่ายขึ้น

6.เอกสารอ้างอิง

- [1] โรงพยาบาลเพชรเวช. (2563). **เหนื่อย หายใจหอบ เท้าบวม สัญญาณเตือนภาวะหัวใจล้มเหลว**. สืบค้นจาก https://www.petcharavejhospital.com/th/Article/article_detail/Heart-Failure.
- [2] จิตตา อริยปรีชากุล. (2562). **ภาวะหัวใจล้มเหลว เกิดขึ้นได้แม้ร่างกายจะดูแข็งแรง**. สืบค้นจาก <https://www.sukumvithospital.com/healthcontent.php?id=3341>.
- [3] สมาคมโปรแกรมเมอร์ไทย. (2563). **การเรียนรู้ของเครื่อง (Machine Learning)** สืบค้นเมื่อ 4 มกราคม 2565. จาก <https://www.thaiprogrammer.org/>.
- [4] เอกสิทธิ์ พัทธวงศ์ศักดิ์. (2558). **การแบ่งข้อมูลเพื่อนำมาทดสอบประสิทธิภาพของโมเดล**. สืบค้น จาก <https://datarockie.com/blog/k-fold-cross-validation/>.
- [5] SKLSONGKIAT. (2559). **ทำไมต้องใช้ K-Nearest Neighbor (K-NN)**. สืบค้นจาก <https://www.sklsongkiat.com/articles/detail/ทำไมต้องใช้-k-nearest-neighbor-k-nn>.
- [6] Kong Ruksiam. (2563). **สรุป Machine Learning(EP.4) —การคำนวณเพื่อนบ้านใกล้เคียงที่สุด (K-nearest Neighbors)**. สืบค้นจาก <https://kongruksiam.medium.com/สรุป-machine-learning-ep-4-เพื่อนบ้านใกล้เคียงที่สุด-k-nearest-neighbors-787665f7c09d>.
- [7] _____. (2563). **การเรียนรู้ของเครื่อง**. สืบค้น จาก <https://th.wikipedia.org/wiki/การเรียนรู้ของเครื่อง>.

- [8] ชาราพิทักษ์วงศ์, จิตราภรณ์. (2560). การใช้ K-Nearest Neighbor Algorithm เพื่อสร้างโมเดลการจำแนกรูปแบบการเรียนรู้สำหรับทำนายผลการจัดระดับมาตรฐานผลิตภัณฑ์ OTOP หัตถกรรม กลุ่มไม้. มหาวิทยาลัยราชภัฏเชียงใหม่.
- [9] พงศกร อีรัมย์. (2558). วิธีการหาค่าเคที่เหมาะสมในการจำแนกแบบเคเนียร์เรสเนเบอร์กับข้อมูลทางการแพทย์. วิทยานิพนธ์นี้เป็นส่วนหนึ่งของการศึกษาตามหลักสูตรปริญญาวิศวกรรมศาสตรมหาบัณฑิต สาขาวิชาวิศวกรรมคอมพิวเตอร์ มหาวิทยาลัยเทคโนโลยีสุรนารี.
- [10] ธวัช ร่มทรัพย์ และ สุรศักดิ์ มั่งสิงห์. (2559). การจำแนกชนิดของพืชด้วยวิธีเพื่อนบ้านใกล้สุดร่วมกับการเลือกตัวแทนที่เหมาะสมด้วยขั้นตอนวิธีเชิงพันธุกรรมโดยใช้คุณลักษณะรูปทรงและพื้นผิวของใบพืช. (ปรัชญาดุษฎีบัณฑิต). สาขาวิชาเทคโนโลยีสารสนเทศ คณะเทคโนโลยีสารสนเทศ มหาวิทยาลัยศรีปทุม.

