

KOMPARASI ALGORITMA DECISION TREE DAN SUPPORT VECTOR MACHINE (SVM) DALAM KLASIFIKASI SERANGAN JANTUNG

Elok Fathiyatul Laili¹, Zakki Alawi², Roihatur Rohmah³, Mula Agung Barata⁴

^{1,4}Teknik Informatika, Universitas Nahdlatul Ulama Sunan Giri

²Sistem Informasi, Universitas Nahdlatul Ulama Sunan Giri

³Sistem Komputer, Universitas Nahdlatul Ulama Sunan Giri

Jl. Ahmad Yani No. 10, Sukorejo, Kabupaten Bojonegoro

e-mail: ¹elokfathiyatul@gmail.com, ²zakki.alawi@unugiri.ac.id,

³roiha.rohmah@unugiri.ac.id, ⁴mula.ab26@gmail.com

Abstract

The heart is one of the most important organs in the human body. According to the WHO, heart attacks are the most common cause of sudden death worldwide, with more than 17.8 million people dying from heart attacks. A heart attack occurs when blood flow to the coronary arteries stops, depriving the heart muscle of oxygen, and causing a heart attack. Detecting a heart attack is very difficult due to the various symptoms. The purpose of this research is to compare the performance of the accuracy values of two algorithms, namely Decision Tree and Support Vector Machine (SVM) in classifying heart attacks. The results of this study show that the Decision Tree algorithm achieves the highest accuracy results compared to the SVM algorithm. The accuracy of the Decision Tree algorithm using a 60:40 ratio data splitting is 98.11% with a negative precision of 98.01% and positive of 98.17% and a negative recall of 97.04% and positive of 98.77%. Meanwhile, the SVM algorithm using data splitting with the same ratio produces an accuracy value of 92.80% with a negative precision of 90.24% and a positive of 94.43% and a negative recall of 91.13% and a positive of 93.85%.

Keyword: Classification, Decision Tree, Heart Attack, Splitting Data, Support Vector Machine

PENDAHULUAN

Jantung adalah salah satu organ yang sangat penting dalam tubuh manusia. Gaya hidup modern yang dijalani masyarakat di kota-kota besar negara maju menjadi salah satu faktor penyebab berbagai masalah kesehatan serius, seperti tekanan darah tinggi, diabetes, serangan jantung hingga komplikasi penyakit lainnya. Hal ini disebabkan oleh kebiasaan makan yang tidak sehat dan kurangnya aktivitas fisik (Purnomo et al., 2020). Serangan jantung adalah penyebab kematian mendadak paling umum di seluruh dunia. Menurut data *Organisasi Kesehatan Dunia* (WHO) lebih dari 17,8 juta orang meninggal karena serangan jantung setiap tahunnya. Di Indonesia angka kematian akibat serangan jantung mencapai 650 ribu setiap tahunnya. Hal ini menyebabkan lebih dari 18 juta orang meninggal setiap tahun karena serangan jantung. Serangan jantung atau juga dikenal sebagai *Infark Miokard Akut (IMA)* terjadi ketika aliran darah arteri koroner terhenti sehingga otot jantung kekurangan oksigen, sehingga menyebabkan infark (Aniamarta et al., 2022). Karena gejalanya dapat beragam dan tidak spesifik, serangan jantung seringkali sulit dideteksi secara dini. Sebagian orang tidak menunjukkan gejala apapun, dan serangan jantung dapat terjadi tanpa gejala, untuk di beberapa kasus ada beberapa gejala serangan jantung yaitu seperti sesak napas, nyeri dada, mual, pusing, muntah serta kelelahan (Gunardi, 2023).

Saat ini perkembangan teknologi informasi mengalami kemajuan yang sangat pesat dan signifikan, dimulai dengan ditemukannya informasi baru dalam big data melalui identifikasi pola-pola tertentu yang sering disebut dengan *data mining* (Qibtiyah & Cahyani, 2023). *Data Mining* adalah proses mengekstraksi dan mengidentifikasi informasi tersembunyi dalam database besar menggunakan metode statistik, kecerdasan buatan dan pembelajaran mesin. Tujuan dari proses ini adalah untuk menemukan atau menggali informasi dan pengetahuan dari data yang ada (Hartawan et al., 2023). Dalam penerapannya *data mining* sering dimanfaatkan untuk menganalisis pola penjualan, memprediksi hasil produksi dan lain sebagainya, namun pada penelitian ini berfokus pada klasifikasi suatu penyakit. Ada beberapa metode dalam *data mining* seperti asosiasi, *clustering*, klasifikasi, serta regresi. Dan salah satu metode yang umum digunakan dalam *data mining* adalah klasifikasi.

Klasifikasi melibatkan pengelompokan objek yang memiliki karakteristik serupa ke dalam beberapa kelas (Azhari et al., 2021). Proses ini disebut juga dengan mempelajari suatu fungsi atau model dari sekumpulan data pelatihan, yang memungkinkan model tersebut memprediksi klasifikasi data pengujian (Septhya et al., 2023). Dalam proses klasifikasi, terdapat berbagai algoritma, termasuk *Decision Tree* dan *SVM*. *Decision Tree* merupakan algoritma yang sangat efektif untuk klasifikasi dan prediksi. Sedangkan *Support Vector Machine (SVM)* merupakan teknik seleksi yang membandingkan sekumpulan parameter standar dengan nilai diskrit, sehingga dapat menghasilkan hasil klasifikasi dapat dicapai dengan tingkat akurasi yang tinggi.

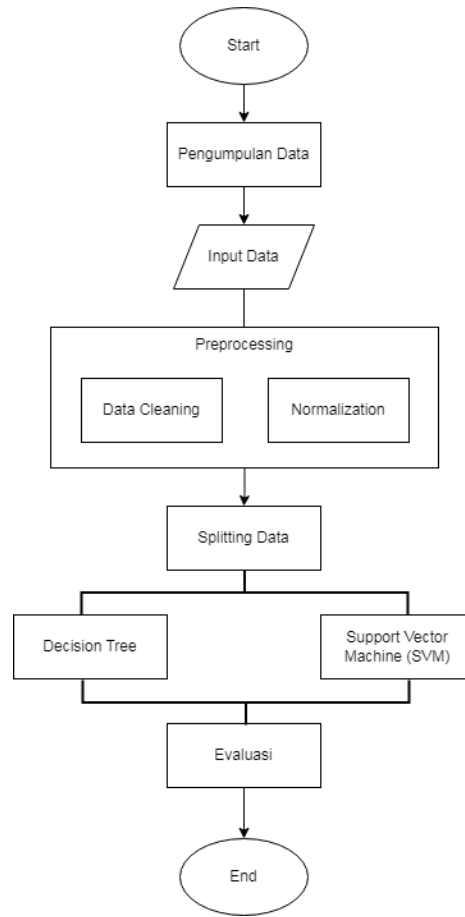
Dalam penelitian yang dilakukan oleh Septhya et al., (2023), dua algoritma dibandingkan, yaitu *Decision Tree* dan *SVM* untuk menentukan algoritma mana yang paling efektif. Selain itu, dalam penelitian ini juga diterapkan *Feature Selection* menggunakan metode *Forward Selection* dengan tujuan untuk meningkatkan akurasi. Hasilnya menunjukkan bahwa *SVM* yang menerapkan *Feature Selection* mencapai akurasi tertinggi, yaitu sebesar 62,3% dengan pembagian data 80:20. Pada penelitian selanjutnya yang dilakukan oleh Munir et al., (2024), juga membahas perbandingan antara dua algoritma pendeteksi kanker payudara yaitu *Decision Tree* dan *Naïve Bayes*. Hasil penelitian ini menunjukkan bahwa *algoritma Decision Tree* lebih unggul dalam mengklasifikasikan kanker payudara, dengan akurasi sebesar 92.04%. Sedangkan algoritma *Naïve Bayes* mencapai akurasi sebesar 91.15%.

Penelitian lebih lanjut yang membandingkan dua algoritma juga dilakukan oleh Hasanah & Nurmalitasari (2023), dengan menerapkan algoritma *Support Vector Machine (SVM)* dan *C4.5*. Hasilnya menunjukkan bahwa algoritma *Support Vector Machine (SVM)* memiliki tingkat akurasi sebesar 87% dan dianggap sebagai algoritma yang paling tepat dan akurat untuk memprediksi penyakit jantung. Penelitian selanjutnya juga dilakukan oleh Ermy Pily et al., (2024), untuk membandingkan performa dari algoritma *K-NN* dan *Naïve Bayes* dalam mengklasifikasikan penyakit diabetes gestasional dengan menerapkan *Feature Selection* dan *splitting data* sebanyak tiga kali, yaitu 60:40, 70:30, dan 80:20. Penelitian ini menggunakan algoritma *K-NN* tanpa *Feature Selection* dan *Naïve Bayes* dengan *Feature Selection* serta rasio *splitting data* 80:20 mendapatkan hasil yang lebih baik.

Berdasarkan penelitian sebelumnya, kedua algoritma yaitu *Decision Tree* dan *SVM* keduanya menunjukkan tingkat akurasi yang tinggi. Dengan demikian, penggunaan algoritma *Decision Tree* dan *SVM* sebagai pilihan untuk mengklasifikasikan serangan jantung dapat menjadi pilihan yang tepat. Oleh karena itu, pada penelitian ini diharapkan dapat memberikan informasi tentang hasil perbandingan nilai akurasi dalam klasifikasi serangan jantung menggunakan algoritma *Decision Tree* dan *Support Vector Machine (SVM)*.

METODE PENELITIAN

Metode penelitian adalah gambaran umum tentang langkah-langkah proses penelitian yang dilakukan untuk mencapai tujuan penelitian. Langkah pertama adalah mengumpulkan data sampai dengan melakukan evaluasi data. Alur tahapan penelitian ini ditunjukkan pada gambar 1.



Gambar 1. Alur Penelitian

1. Pengumpulan Data

Pengumpulan data merupakan tahap pertama dari penelitian ini. Data dapat diambil dari berbagai sumber, salah satunya dengan mencari dataset di internet dimana dataset tersebut merupakan data publik (Ermy Pily et al., 2024).

Data penelitian ini berdasarkan data *Heart Attack Classification* yang tersedia di platform *Kaggle* di alamat web yaitu: <https://www.kaggle.com/datasets/thxogg/heart-attack-classification-training-dataset>. Dataset ini berjumlah 1319 data dengan 8 fitur faktor penyebab serangan jantung dan 1 kelas klasifikasi serangan jantung. Adapun 8 fitur yang menjadi faktor penyebab serangan jantung yaitu: *age*, *gender*, *impulse*, *pressurehigh*, *pressurelow*, *glucose*, *kcm*, dan *troponin*. Data ditunjukkan pada Tabel 1.

Tabel 1. Fitur Dataset

No	Fitur	
	Atribut	Deskripsi
1	<i>age</i>	Usia
2	<i>gender</i>	jenis kelamin (0 = perempuan, 1 = laki-laki)
3	<i>impulse</i>	sinyal listrik yang dihasilkan oleh nodus sinus
4	<i>pressurehigh</i>	tekanan darah tinggi
5	<i>pressurelow</i>	tekanan darah rendah
6	<i>glucose</i>	kadar gula darah dalam tubuh
7	<i>kcm</i>	kerusakan pada otot jantung
8	<i>troponin</i>	peningkatan kadar troponin dalam darah

Tabel 2. Kelas Dataset

No	Fitur	
	Kelas	Deskripsi
1	<i>class</i>	menunjukkan kelas kategori (<i>negative</i> = tidak terkena serangan jantung, <i>positive</i> = terkena serangan jantung)

2. Preprocessing Data

Preprocessing data adalah tahapan untuk melakukan pembersihan pada suatu data yang akan digunakan. Tahap ini sangat penting karena data yang digunakan dalam proses data mining tidak selalu berada pada tempat yang tepat untuk diproses. Hal ini karena data tersebut mungkin mengandung beberapa permasalahan yang dapat mempengaruhi hasil proses *data mining*, seperti terdapat nilai yang hilang, *outliner*, atau format data yang tidak sesuai. Oleh karena itu, tahap *preprocessing* harus dilakukan untuk mengatasi permasalahan tersebut (Dwifabri et al., 2021). Berikut adalah beberapa tahap *preprocessing* data yang dilakukan dalam penelitian ini:

- 1) *Data Cleaning* atau juga dikenal sebagai pemebersihan data, adalah tahap pertama dalam *preprocessing* data. Tujuannya untuk membersihkan data duplikat dan data kosong atau tidak relevan agar dapat diolah dan dianalisis dengan baik.
- 2) *Normalization Data* adalah proses transformasi atribut numerik dengan skala ke format yang lebih sederhana, seperti 0-1. Data yang sudah dilakukan proses akan berkisar antara 0-1. Ada beberapa metode yang dapat digunakan untuk menormalisasikan data, seperti persamaan 1 dan 2 sebagai berikut :
 - *Min-Max Normalization*, yaitu data ditransformasikan secara linier dari satu nilai kenilai baru.

$$v' = \frac{v - \min_A}{\max_A - \min_A} (\text{new_max}_A - \text{new_min}_A)$$

(1)

- *Z-Score Normalization*, yaitu untuk melakukan input dengan menggunakan rata-rata standar deviasi.

$$v' = \frac{v - \bar{A}}{\sigma_A}$$

(2)

Tabel 3. Dataset Serangan Jantung setelah *Preprocessing*

age	gender	impluse	pressurehight	Pressurelow	glucose	kcm	troponin	class
62	1	66	160	83	160.0	1.80	0.012	0
21	1	96	98	46	296.0	6.75	1.06	1
55	1	64	160	77	270.0	1.99	0.003	0
64	1	70	120	55	270.0	13.87	0.122	1
...	...	...	...	...	...	...	...	...
45	1	85	168	104	96.0	1.24	4.25	1
54	1	58	117	68	443.0	5.8	0.359	1
51	1	94	157	79	134.0	50.89	1.77	1

3. Splitting Data

Splitting Data adalah proses pembagian data menjadi dua bagian, yaitu data *training* dan data *testing* (Eina et al., 2024). Pemisahan ini dilakukan dengan rasio tertentu, untuk memastikan model dapat diterapkan pada data baru. Data *training* digunakan untuk mengevaluasi keberhasilan atau kegagalan penelitian, sementara data

testing digunakan untuk melatih model (Putri et al., 2024). Dari penelitian yang dilakukan oleh (Ermy Pily et al., 2024), proses dalam penelitian ini dengan *splitting data* tiga kali dengan rasio 60:40, 70:30, 80:20. Sehingga dalam penelitian ini juga menggunakan proses tiga kali *splitting data* dengan rasio yang sama, namun terdapat perbedaan pada penerapan algoritma. Data yang sudah dibagi akan dilakukan data *testing* menggunakan algoritma *Decision Tree* dan *SVM*.

4. Klasifikasi

Klasifikasi merupakan salah satu metode *data mining* yang berfokus pada pengelompokan data berdasarkan kriteria tertentu melalui analisis data yang ada (Rahayu et al., 2023). Prinsip klasifikasi adalah memprediksi label kelas kategorikal dari data yang diberikan, sehingga data tersebut dapat dikelompokkan ke dalam salah satu kelas yang telah ditentukan (Barata et al., 2023). Klasifikasi dapat dilakukan dengan berbagai macam metode, antara lain penggunaan *machine learning* dengan algoritma seperti *Decision Tree*, *Support Vector Machine* (SVM), *Naïve Bayes*, dan *Random Forest* (Afifah et al., 2023).

5. Decision Tree

Decision Tree merupakan algoritma klasifikasi yang sangat populer. Algoritma ini menggunakan pendekatan pohon keputusan dengan memilih suatu atribut sebagai akar, kemudian membuat cabang untuk setiap nilai pada akar, dan mengulangi proses tersebut untuk di setiap cabang hingga kasus pada cabang tersebut berada pada kelas yang sama. Langkah-langkah dalam proses membuat pohon keputusan yaitu (Suriani, 2023) :

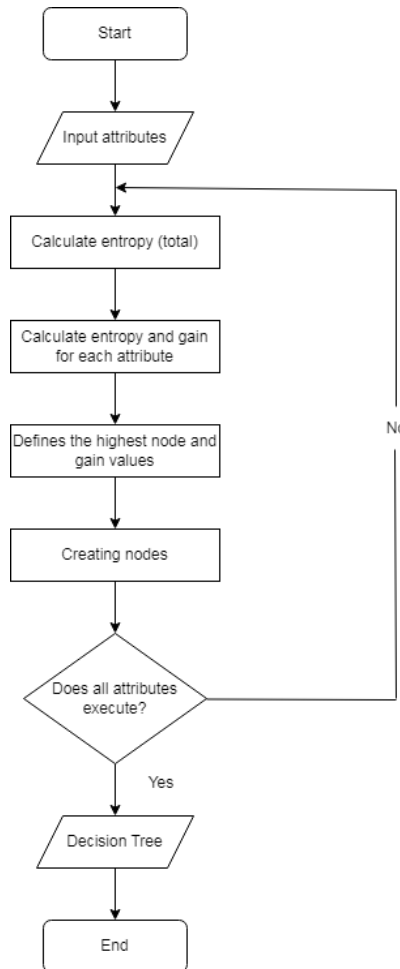
- 1) Menyiapkan data *training*
- 2) Menentukan akar pohon keputusan, akar ini dipilih dengan menghitung gain masing-masing atribut.
- 3) Menghitung nilai *gain*, nilai *gain* gunakan untuk menentukan atribut mana yang akan diuji pada setiap node di pohon keputusan, dan pilih atribut dengan informasi *gain* tertinggi sebagai atribut pengujian untuk *node* tersebut (Hasanah & Nurmalitasari, 2023). Perhitungan gain dapat dilakukan dengan menggunakan rumus yang terdapat dalam persamaan 3 :

$$Gain(S, A) = Entropy(S) - \sum_{i=1}^n \frac{|S_i|}{|S|} \times Entropy(S_i) \quad (3)$$

- 4) Ulangi prosedur kedua hingga setiap cabang terpenuhi. Adapun untuk menghitung nilai *entropy* dapat dilihat pada persamaan 4 di bawah ini :

$$Entropy(S) = \sum_{i=1}^n - \pi \times \log_2 \pi \quad (4)$$

- 5) Proses partisipasi akan berakhir jika semua cabang *node* mendapat kelas yang sama



Gambar 2. Alur *Decision Tree*

6. Support Vector Machine (SVM)

SVM merupakan termasuk dalam algoritma dalam *Supervised Learning* yang digunakan untuk menganalisis data dan mengidentifikasi pola dalam klasifikasi data. Untuk membagi data ke dalam dua kelas, *SVM* mencari *hyperlane* (pemisah) terbaik dan memaksimalkan jarak antara dua kelas tersebut (Suryani & Mustakim, 2022). Namun, algoritma ini juga memiliki kekurangan terkait pemilihan parameter atau fitur yang tepat, yang dapat mempengaruhi hasil akurasi klasifikasi (Azhari et al., 2021). Algoritma *SVM* memiliki beberapa jenis kernel berdasarkan fungsinya dalam mengatasi masalah pada algoritma ini, yaitu dapat dilihat pada persamaan 5, 6 dan 7 (Riyantoni et al., 2023) :

- 1) Kernel Linier

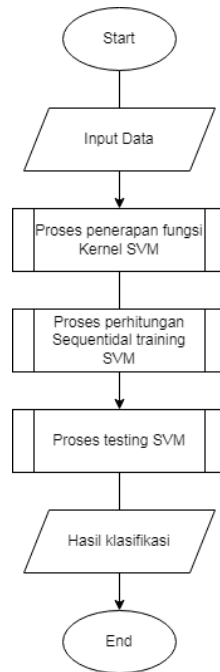
$$K(x_i, x) = x_i^T x \quad (5)$$

- 2) Kernel Polynomial

$$K(x_i, x) = (\gamma(x_i^T x) + r)^p \quad (6)$$

- 3) Kernel Radial Basis Function (RBF)

$$K(x_i, x) = \exp(-\gamma \|x_i - x\|^2 + C) \quad (7)$$



Gambar 3. Alur Support Vector Machine (SVM)

7. Google Colab

Google Colab (Colaboratory) merupakan layanan gratis dari Google yang digunakan untuk menulis dan menjalankan kode *Python* melalui *browser*. Platform ini berbasis cloud yang memungkinkan penggunaanya menjalankan serta berbagi notebook tanpa harus menginstal atau mengkonfigurasinya (Guntara, 2023). *Google Colab* sangat populer di berbagai kalangan, karena platform ini menawarkan banyak fitur yang sangat berguna baik untuk pemula maupun profesional terutama dalam bidang data *science*, *machine learning*, dan pengolahan data secara umum.

8. Evaluasi

Evaluasi dilakukan setelah proses pengujian algoritma selesai. Tujuan dari evaluasi dan perbandingan algoritma yaitu untuk membuktikan bahwa model benar-benar mempresentasikan sesuai pemodelan yang dilakukan serta mengukur performa dari dua algoritma yang digunakan.

Pada penelitian ini *Confusion Matrix* digunakan untuk menilai akurasi hasil dari kedua algoritma. *Confusion Matrix* mampu memberikan informasi mengenai hasil klasifikasi yang sudah dilakukan oleh *Machine Learning* (Apriyani & Kurniati, 2020).

Tabel 4.Confusion Matix

		Kelas Prediksi	
		True	False
Kelas Aktual	True	True Positive (TP)	False Negative (FN)
	False	False Positive (FP)	True Negative (TN)

Ada tiga nilai metriks yang digunakan untuk mengevaluasi kinerja sistem klasifikasi yang dikembangkan, yaitu *Precision*, *Recall*, dan *Accuracy* (Azhari et al., 2021). Nilai-nilai *Precision*, *Recall*, dan *Accuracy* dapat dihitung dengan menggunakan rumus yang terdapat dalam persamaan 8, 9 dan 10 berikut :

$$Precision = \frac{(TP + TN)}{TP + FP} \times 100\% \quad (8)$$

$$Recall = \frac{TP}{TP + FN} \times 100\% \quad (9)$$

$$Accuracy = \frac{(TP + TN)}{TP + TN + FP + FN} \times 100\% \quad (10)$$

HASIL DAN PEMBAHASAN

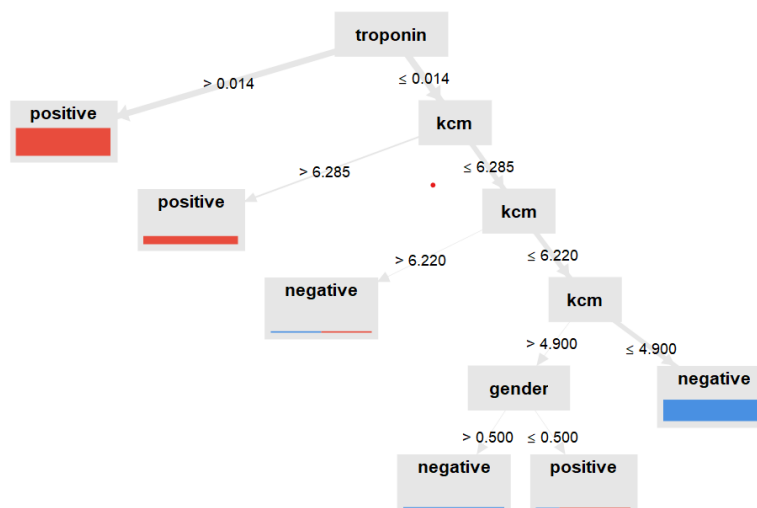
Berdasarkan hasil pengujian yang telah dilakukan dengan menggunakan kedua algoritma menunjukkan bahwa masing-masing algoritma menghasilkan nilai akurasi yang berbeda pada kasus serangan jantung. Dalam pengujian ini dataset yang digunakan diperoleh dari data publik yang tersedia pada platform *Kaggle*, yaitu dengan menggunakan 1319 dataset dengan 8 *feature* dan 1 kelas klasifikasi. Pada tahapan *preprocessing* data dilakukan data *cleaning* dan *normalization data*.

Setelah melakukan *preprocessing* data, langkah selanjutnya adalah memodelkan algoritma *Decision tree* dan *SVM* dengan melakukan tiga kali *splitting data* menggunakan rasio perbandingan 60:40, 70:30, dan 80:20. Proses perhitungan dan pemodelan dilakukan menggunakan platform *Google Colab*, yang memungkinkan pengolahan data secara efisien dan mendukung hasil analisis akurasi dalam pengolahan data. Melalui *Google Colab*, hasil akurasi algoritma *Decision Tree* dan *SVM* dapat dibandingkan secara optimal. Berikut hasil nilai akurasi pemodelan algoritma yang sudah dibangun dengan menggunakan *Google Colab* pada tabel 5.

Tabel 5. Perbandingan Akurasi Algoritma

Algoritma	Splitting Data	Akurasi	Precision		Recall	
			Negative(0)	Positive(1)	Negative(0)	Positive(1)
Decision Tree	60:40	98.11%	98.01%	98.17%	97.04%	98.77%
	70:30	97.73%	98.03%	97.54%	96.13%	98.76%
	80:20	97.35%	97.00%	97.56%	96.04%	98.16%
SVM	60:40	92.80%	90.24%	94.43%	91.13%	93.85%
	70:30	91.41%	89.54%	92.59%	88.39%	93.36%
	80:20	92.80%	90.24%	94.43%	91.13%	93.85%

Pada tabel 5 menunjukkan hasil perbandingan nilai akurasi yang terbaik untuk kedua algoritma. Dalam pengujian yang telah dilakukan menggunakan *Platform Google Colab*, algoritma *Decision Tree* mendapatkan nilai akurasi tertinggi yaitu sebesar 98.11% dengan *splitting* data 60:40. Sedangkan untuk algoritma *SVM* mendapatkan nilai akurasi tertinggi sebesar 92.80% dengan *splitting* data yang sama.



Gambar 4. Pohon Keputusan Algoritma *Decision Tree*

KESIMPULAN

Berdasarkan penelitian yang telah dilakukan pada dataset *Heart Attack Classification* yang diambil dari data publik melalui platform *Kaggle* yang berjumlah 1319 data, maka dapat disimpulkan bahwa implementasi dan perbandingan algoritma *Decision Tree* dan *SVM* berhasil dilakukan. Seluruh proses analisis hingga evaluasi dilakukan menggunakan platform *Google Colab*, yang mendukung efisiensi dalam pengolahan data. Dengan membandingkan nilai akurasi terlihat bahwa algoritma *Decision Tree* memberikan hasil performa yang lebih akurat dibandingkan dengan algoritma *Support Vector Machine (SVM)*. Algoritma *Decision Tree* yang menggunakan *splitting data* dengan rasio perbandingan 60:40 menghasilkan nilai akurasi sebesar 98.11% dengan *precision negative* sebesar 98.01% dan *positive* sebesar 98.17% serta *recall negative* sebesar 97.04% dan *positive* sebesar 98.77%. Dan, algoritma *SVM* menggunakan *splitting data* dengan rasio yang sama menghasilkan nilai akurasi sebesar 92.80% dengan *precision negative* sebesar 90.24% dan *positive* sebesar 94.43% serta *recall negative* sebesar 91.13% dan *positive* sebesar 93.85%.

SARAN

Berdasarkan dari hasil penelitian yang telah dilakukan, diperlukannya penelitian lebih lanjut untuk meningkatkan pemahaman mengenai analisis. Salah satu cara untuk melakukannya adalah dengan menguji beberapa algoritma pada kumpulan data yang lebih besar, memiliki fitur yang lebih beragam, mengeksplorasi algoritma lain, dan menambahkan fitur yang dapat meningkatkan nilai akurasi.

DAFTAR PUSTAKA

- Afifah, T. A., Martiansah, R., Alviyoni, M., & Arifin, A. (2023). Perbandingan Algoritma Klasifikasi C4.5 Dan Naive Bayes untuk Memprediksi Gagal Jantung. *SENTIMAS: Seminar Nasional Penelitian Dan Pengabdian Masyarakat*, 78–85.
- Aniamarta, T., Salsabilla Huda, A., & Lizariani Aqsha, F. (2022). Causes and Treatments of Heart Attack. *Biologica Samudra*, 4(1), 22–31. <https://doi.org/10.33059/jbs.v4i1.3925>
- Apriyani, H., & Kurniati, K. (2020). Perbandingan Metode Naive Bayes Dan Support Vector Machine Dalam Klasifikasi Penyakit Diabetes Melitus. *Journal of Information Technology Ampera*, 1(3), 133–143. <https://doi.org/10.51519/journalita.volume1.issue3.year2020.page133-143>
- Azhari, M., Situmorang, Z., & Rosnelly, R. (2021). Perbandingan Akurasi, Recall, dan Presisi Klasifikasi pada Algoritma C4.5, Random Forest, SVM dan Naive Bayes. *Jurnal Media Informatika Budidarma*, 5(2), 640. <https://doi.org/10.30865/mib.v5i2.2937>
- Barata, M. A., Edi Noersasongko, Purwanto, & Moch Arief Soeleman. (2023). Improving the Accuracy of C4.5 Algorithm with Chi-Square Method on Pure Tea Classification Using Electronic Nose. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 7(2), 226–235. <https://doi.org/10.29207/resti.v7i2.4687>
- Dini Serangan Jantung Saat Aktivitas Olahraga, D., & Gunardi, S. (2023). Early Detection of Heart Attack During Sports Activities. *Jurnal Hasil Kegiatan Pengabdian Masyarakat Indonesia*, 1(3), 257–265.
- Dwifabri, M., Irvan, M., Imam, A., & Adiwijaya. (2021). *JURNAL RESTI Perbandingan Support Vector Machine dan Modified Balanced Random*. 1(10), 393–399.
- Eina, M. F., Chrisnanto, Y. H., & Melina, M. (2024). Klasifikasi Telemarketing Menggunakan Naive Bayes Classification Dan Wrapper Sequential Feature Selection. *INTECOMS: Journal of Information Technology and Computer Science*, 7(4), 1189–1198. <https://doi.org/10.31539/intecom.v7i4.10846>
- Ermy Pily, A. K., Oktavianda, Aprilia, F., Rahmadden, & Efrizoni, L. (2024). Komparasi Algoritma K-Nearest Neighbors dan Naive Bayes dalam Klasifikasi Penyakit Diabetes Gestasional. *Indonesian Journal of Computer Science*, 13(1), 1195–1209. <https://doi.org/10.33022/ijcs.v13i1.3714>
- Guntara, R. G. (2023). *Deteksi Atap Bangunan Berbasis Citra Udara Menggunakan Google Colab*

- dan Algoritma Deep Learning YOLOv7. 2(1), 9–18.
- Hartawan, M. S., Erkamim, M., Rachmawati, S., Santi, N. C., Legito, L., & Sepriano, S. (2023). Penerapan Algoritma Supervised Learning untuk Klasifikasi Program Keluarga Harapan. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 3(2), 83–91. <https://doi.org/10.57152/malcom.v3i2.873>
- Hasanah, H., & Nurmalitasari. (2023). Perbandingan Tingkat Akurasi Algoritma Support Vector Machines (SVM) dan C45 dalam Prediksi Penyakit Jantung. *Prosiding Seminar Nasional Teknologi Dan Sains*, 2, 13–18.
- Munir, A. S., Saputra, A. B., Aziz, A., & Barata, M. A. (2024). Perbandingan Akurasi Algoritma Naive Bayes dan Algoritma Decision Tree dalam Pengklasifikasian Penyakit Kanker Payudara. *Jurnal Ilmiah Informatika Global*, 15(1), 23–29. <https://doi.org/10.36982/jiig.v15i1.3578>
- Purnomo, A., Barata, M. A., Soeleman, M. A., & Alzami, F. (2020). Adding feature selection on Naïve Bayes to increase accuracy on classification heart attack disease. *Journal of Physics: Conference Series*, 1511(1). <https://doi.org/10.1088/1742-6596/1511/1/012001>
- Putri, A. D., Sholekhah, F., & Dadynata, E. (2024). *The Application of C4 . 5 Decision Tree Algorithm for Predicting the Survival Rate of Thyroid Cancer Patients Penerapan Algoritma Decesion Tree C4 . 5 untuk Memprediksi Tingkat Kelangsungan Hidup Pasien Kanker Tiroid*. 4(October), 1485–1495.
- Qibtiyah, M., & Cahyani, N. (2023). *Prediksi Tingkat Kelahiran Bayi Di Kabupaten Bojonegoro Dengan Menggunakan Algoritma Naive Bayes*. 1–7.
- Rahayu, D. S., Afifah, J., & Intan, S. (2023). Classification of Diabetes Mellitus Using C4 . 5 Algorithm , Support Vector Machine (SVM) and Linear Regression Klasifikasi Penyakit Diabetes Melitus Menggunakan Algoritma C4 . 5 , Support Vector Machine (SVM) dan Regresi Linear. *SENTIMAS: Seminar Nasional Penelitian Dan Pengabdian Masyarakat*, 1(1 SE-), 56–63.
- Riyantoni, N. H., Bahreisy, M. F., Hakim, I., & Rolliawati, D. (2023). Komparasi Support Vector Machine (SVM) Dan Autoregressive Integrated Moving Average (Arima) Pada Peramalan Hujan Di Daerah Albury, Australia. *Jurnal Sistem Informasi Dan Informatika (Simika)*, 6(1), 59–68. <https://doi.org/10.47080/simika.v6i1.2412>
- Septhya, D., Rahayu, K., Rabbani, S., Fitria, V., Rahmadden, R., Irawan, Y., & Hayami, R. (2023). Implementasi Algoritma Decision Tree dan Support Vector Machine untuk Klasifikasi Penyakit Kanker Paru. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 3(1), 15–19. <https://doi.org/10.57152/malcom.v3i1.591>
- Suriani, U. (2023). Penerapan Data Mining untuk Memprediksi Tingkat Kelulusan Mahasiswa Menggunakan Algoritma Decision Tree C4.5. *Journalcisa*, 3(2), 55–66.
- Suryani, S., & Mustakim, M. (2022). Estimasi Keberhasilan Siswa dalam Pemodelan Data Berbasis Learning Menggunakan Algoritma Support Vector Machine. *Bulletin of Informatics and Data Science*, 1(2), 81. <https://doi.org/10.61944/bids.v1i2.36>