

1214 KLASIFIKASI DIAGNOSA PENYAKIT TIROID MENGGUNAKAN METODE RANDOM FOREST

By Raja Darmawan

KLASIFIKASI DIAGNOSA PENYAKIT TIROID MENGGUNAKAN METODE RANDOM FOREST

Raja Darmawan

Abstract

The thyroid is a vital gland in the human neck, regulating metabolism through its hormones. Hormonal disorders in the thyroid can significantly affect health. Data mining techniques, such as the random forest algorithm, are used to analyze thyroid disease data. Previous research has used methods such as Decision Tree and Support Vector Machine with high accuracy results. This study aims to apply the random forest method in the classification of thyroid disease diagnoses. Research questions include factors that affect accuracy, the effect of parameter changes on the model, and optimal data sharing. The results show that parameters such as the number of decision tree maximum depth, and minimum number of samples can affect the accuracy of the model. The evaluation showed that the highest accuracy was obtained in the first with a data split of 80/20 with an accuracy result of 99%. This study concludes that the random forest method is effective in improving the accuracy of thyroid disease diagnosis and the importance of parameter adjustment for optimal results.

Keywords : *Thyroid Disease, Classification, Random Forest*

Abstrak

Tiroid merupakan kelenjar vital di leher manusia, mengatur metabolisme melalui hormonnya. Gangguan hormonal pada tiroid dapat mempengaruhi kesehatan secara signifikan. Teknik data mining, seperti algoritma random forest, digunakan untuk menganalisis data penyakit tiroid. Penelitian sebelumnya telah menggunakan metode seperti Decision Tree dan Support Vector Machine dengan hasil akurasi tinggi. Penelitian ini bertujuan untuk menerapkan metode random forest dalam klasifikasi diagnosa penyakit tiroid. Pertanyaan penelitian termasuk faktor-faktor yang memengaruhi akurasi, pengaruh perubahan parameter terhadap model, dan pembagian data yang optimal. Hasil penelitian menunjukkan bahwa parameter seperti jumlah pohon keputusan, kedalaman maksimum, dan jumlah minimum sampel dapat memengaruhi akurasi model. Evaluasi menunjukkan akurasi tertinggi diperoleh pada pengujian pertama dengan pembagian data 80/20 dengan hasil akurasi diperoleh 99%. Penelitian ini menyimpulkan bahwa metode random forest efektif dalam meningkatkan akurasi diagnosa penyakit tiroid dan pentingnya penyesuaian parameter untuk hasil yang optimal.

Kata kunci : *Penyakit Tiroid, Klasifikasi, Random Forest*

1. PENDAHULUAN

Bagian tubuh manusia yang paling penting adalah tiroid, yang terletak di bawah jakun pada leher. Kelenjar ini, yang memiliki bentuk seperti kupu-kupu, menghasilkan hormon tiroid dan bertanggung jawab untuk menjaga metabolisme [1]. Sangat penting untuk mengamati dan mengukur berbagai gangguan hormonal pada manusia. Faktor penyebab tiroid adalah fleksibilitas hormon manusia yang terlalu tinggi atau rendah. Perawatan kesehatan penting untuk

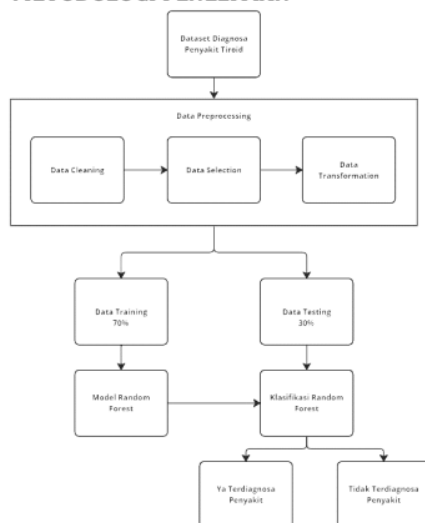
menyeimbangkan variasi hormon yang terlalu tinggi atau rendah. Faktor risiko untuk gangguan hormon cukup banyak, sehingga dokter membutuhkan perhatian yang lebih besar untuk memantau dan menemukan diagnosis yang tepat pada waktu yang tepat [2]. Karena gejala gangguan fungsi tiroid sangat mirip dengan keluhan yang disebabkan oleh gaya hidup modern, gangguan fungsi tiroid sering kali sulit diidentifikasi. Akibatnya, pasien sering mengabaikan masalahnya dan tidak memeriksanya ke dokter [3].

Oleh karena itu, Teknik data mining diperlukan untuk menganalisa data penyakit ini [4]. Dalam Teknik Data mining terdapat berbagai algoritma klasifikasi seperti *naive bayes*, *decision tree*, *neural network*, *k-random forest*, dan sebagainya. untuk menggambarkan dan membedakan kelas data atau konsep. [5]. Teknik data mining yang digunakan adalah algoritma klasifikasi random forest. Random forest adalah konsep pembelajaran kelompok yang menggunakan hasil dari beberapa pohon keputusan untuk meningkatkan akurasi. Metode ini bekerja dengan membangun pohon keputusan secara acak dari kumpulan data. Teknik ini dapat menghasilkan tingkat akurasi yang tinggi dan menghindari masalah overfitting, tetapi juga memerlukan waktu pelatihan yang lebih lama [6].

Berdasarkan penelitian terdahulu melakukan penelitian untuk penyakit tiroid dengan menerapkan metode *Decision Tree* dan *Support Vector Machine*, didapatkan hasil nilai akurasi sebesar 99.89% [7]. Penelitian tambahan tentang pengujian dengan menggunakan dataset penyakit tiroid sebanyak 500 data, algoritma *Naive Bayes* memiliki akurasi sebesar 63% pada data *testing* dan 64% pada data *training*, sedangkan algoritma *K-Nearest Neighbors* (KNN) memiliki akurasi sebesar 63% [8].

Berdasarkan latar belakang yang telah disebutkan, penelitian ini bertujuan untuk mengetahui bagaimana penerapan metode random forest dalam proses klasifikasi diagnosa penyakit tiroid dengan hasil performa terbaik serta membuat model klasifikasi yang sederhana dan memiliki performa baik.

2. METODOLOGI PENELITIAN



Gambar 1. Alur Penelitian

Data mining merupakan teknik yang untuk mengekstrak informasi dan pengetahuan penting dari kumpulan data yang sangat besar. Informasi dan pengetahuan ini dapat digunakan untuk berbagai tujuan, seperti riset ilmiah, analisis bisnis, dan pengambilan Keputusan [9]. Sederhananya, data mining adalah penemuan pola pengolahan data yang bermanfaat [10]. Data mining merupakan proses mengekstrak informasi dari data dalam jumlah besar. Informasi tersebut dapat berupa fakta, kesimpulan, pola, atau anomali. Data mining memiliki lima fungsi utama, yaitu [11]:

- Klasifikasi adalah proses memprediksi kelas data baru berdasarkan data yang telah dilatih sebelumnya.
- *Clustering* adalah proses membagi data ke dalam kelompok-kelompok berdasarkan atributnya.
- *Association* adalah proses mengidentifikasi hubungan antara dua atau lebih variabel.
- *Sequencing* adalah proses mengidentifikasi urutan kejadian atau peristiwa.
- *Forecasting* adalah proses memprediksi nilai pada masa yang akan datang berdasarkan data masa lalu.

2.1. Pengumpulan Data

Pada penelitian ini data yang digunakan didapatkan dari Kaggle. Data ini menangkap berbagai atribut terkait kondisi tiroid untuk diagnosis medis. Ini mencakup informasi demografis seperti usia dan jenis kelamin. Fitur riwayat kesehatan mencakup asupan tiroksin, obat antitiroid, dan operasi sebelumnya. Status kesehatan pasien terkini mengenai penyakit, kehamilan, adanya penyakit gondok, tumor, atau kondisi hipofisis dicatat. Selain itu, ia mencatat apakah pasien mencurigai hipotiroidisme atau hipertiroidisme. Hasil laboratorium seperti kadar TSH, T3, TT4, T4U, FTI disertakan jika diukur. Klasifikasi biner menunjukkan ada tidaknya hipertiroidisme. Sumber rujukan juga ditunjukkan Dataset berisi atribut-atribut.

2.2. Preprocessing

Preprocessing adalah tahap penting dalam analisis data yang dilakukan untuk mengubah data mentah dari proses pengambilan data menjadi format yang lebih efisien untuk digunakan dalam proses klasifikasi [17]. Proses ini dilakukan untuk memeriksa dan memperbaiki kesalahan pada dataset. Teknik yang digunakan pada penelitian ini yaitu *cleaning*, *selection*, dan *transformation*. *Cleaning* diperlukan untuk mengatasi data yang hilang menggunakan hasil

rata-rata data. Lalu proses *selection* dilakukan untuk menunjukan kontribusi relative setiap fitur dalam membuat keputusan. Semakin tinggi nilai penting suatu fitur, semakin besar pengaruh fitur tersebut dalam model. Selain itu, proses *transformation* digunakan untuk merubah data dengan tipe *categorical* menjadi *numerik* yaitu 0 dan 1.

2.3. Split Data

Langkah selanjutnya adalah memisahkan data. Pemisahan data adalah proses membagi data menjadi dua bagian yaitu data uji dan data latih [16]. Setelah melewati proses *processing*, kemudian dilakukan split data untuk data *training* dan data *testing*. Pada bagian data pada penelitian ini menggunakan 70/30. 70% data *training* dan 30% data *testing*.

2.4. Klasifikasi Random Forest

Klasifikasi adalah proses pembelajaran mesin yang menggunakan fitur-fitur dari suatu objek untuk memprediksi kelas objek tersebut. Metode ini bekerja dengan membangun model yang dapat mengidentifikasi ciri-ciri yang membedakan satu kelas dengan kelas lainnya. Selanjutnya, model ini dapat digunakan untuk memprediksi kelas objek baru yang label kelasnya belum diketahui [10]. Klasifikasi adalah proses memprediksi kelas data baru dengan menggunakan data yang telah dilatih sebelumnya. Data yang telah dilatih disebut data *training*, sedangkan data baru disebut data *testing*. Kemungkinan kelas data baru ditentukan berdasarkan label atau atribut tujuan dari data *training* [12].

Random Forest merupakan kumpulan pohon keputusan yang digunakan untuk mengklasifikasi data ke suatu kelas. Ini adalah salah satu teknik yang dapat meningkatkan akurasi hasil dalam menciptakan atribut untuk setiap node yang dilakukan secara acak. Pohon keputusan dibangun dengan menentukan node akar, lalu cabang-cabang pohon akan bercabang lagi hingga mencapai node daun. Hasil akhir dari pohon keputusan adalah kelas yang diprediksi untuk data baru [13]. Metode Random Forest digunakan untuk melakukan klasifikasi dan regresi. Terdapat tiga poin kunci: (1) pembangunan pohon prediksi dengan melakukan bootstrap sampling; (2) Setiap pohon keputusan melaksanakan memanfaatkan variabel prediktor secara acak; (3) hasil dari setiap pohon keputusan digabungkan oleh Random Forest menggunakan mayoritas suara untuk klasifikasi atau rata-rata untuk regresi. Metode ini menggabungkan hasil dari berbagai pohon keputusan untuk meningkatkan akurasi dan kinerja model [14].

Metode Random Forest memiliki dua parameter utama, yaitu jumlah pohon keputusan (*tree*) dan jumlah variabel independen yang digunakan untuk proses pencabangan. Untuk menentukan nilai parameter *mtry*, disarankan menggunakan persamaan berikut [15]:

$$mtry_1 = \frac{1}{2} \sqrt{p} \quad (1)$$

$$mtry_2 = \sqrt{p} \quad (2)$$

$$mtry_3 = 2 \times \sqrt{p} \quad (3)$$

mtry = jumlah variabel independen yang signifikan untuk setiap split
p = jumlah variabel independent

Pada tahap ini dilakukan pembuatan model random forest untuk klasifikasi diagnose penyakit tiroid. Random forest ini memiliki beberapa parameter yang diujikan untuk mengoptimalkan model. Beberapa parameter tersebut adalah *n_estimator* (jumlah pohon keputusan), *max_depth* (kedalaman maksimum pohon keputusan), *min_sample_leaf* (jumlah minimum sampel daun), dan *random state* (nilai inisialisasi).

2.5. Evaluasi

Setelah itu dilakukan proses evaluasi terhadap performa model random forest. Evaluasi ini melibatkan penggunaan *matrices* seperti confusion matrix untuk mengukur metrik-metrik penting seperti akurasi, presisi, dan recall. Evaluasi ini membantu dalam memahami seberapa baik model tersebut dalam mengklasifikasi data. Confusion matrix adalah tabel yang digunakan untuk mengukur akurasi dan kinerja algoritma klasifikasi. Tabel ini menunjukkan berapa banyak data *testing* yang diprediksi dengan benar dan salah oleh algoritma tersebut. Confusion matrix dapat digunakan untuk menilai kinerja algoritma klasifikasi dalam menyelesaikan masalah klasifikasi [13]. Confusion matrix merupakan sebuah metode yang menghitung *accuracy*, *recall*, dan *precision* melalui tabel matriks [15]:

TABEL I. CONFUSION MATRIX

		True Class	
		Positive	Negative
Predicted Class	Positive	True positive (TP)	False Negative (FN)
	Negative	False Positive (FP)	True Negative (TN)

$$Accuracy = \frac{TP+TN}{TP+TN+FN+FP} \quad (1)$$

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

$$Recall = \frac{TP}{TP+FN} \quad (1)$$

3. HASIL DAN PEMBAHASAN

3.1. Pengumpulan Data

TABEL II. DATASET

No	Atribut	Keterangan
1	Age	Usia pasien
2	Sex	Jenis kelamin pasien
3	Thyroxine	Apakah pasien dalam pengobatan tiroksin
4	Queryonthyroxine	Apakah pasien dalam pengobatan tiroksin
5	onantithyroidmedication	Apakah pasien dalam pengobatan antitiroid
6	Sick	Apakah pasien sedang sakit
7	Pregnant	Apakah pasien sedang hamil
8	thyroidsugery	Apakah pasien telah menjalani operasi tiroid
9	I131treatment	Apakah pasien sedang menjalani pengobatan I131
10	Queryhypothyroid	Apakah pasien yakin menderita hipotiroid
11	Queryhyperthyroid	Apakah pasien yakin menderita hipertiroid
12	Lithium	Apakah pasien menggunakan lithium
13	Goitre	Apakah pasien memiliki gondok
14	Tumor	Apakah pasien

		memiliki tumor
15	Hypopituitary	Apakah pasien memiliki kelenjar hipofisis
16	Psych	Apakah pasien memiliki masalah kejiwaan
17	TSHmeasured	Apakah TSH diukur dalam darah
18	TSH	Kadar TSH dalam darah
19	T3measured	Apakah T3 diukur dalam darah
20	T3	Kadar T3 dalam darah
21	TT4measured	Apakah TT4 diukur dalam darah
22	T4	Kadar T4 dalam darah
23	T4Umeasured	Apakah T4U diukur dalam darah
24	T4U	Kadar T4U dalam darah
25	FTImeasured	Apakah FTI diukur dalam darah
26	FTI	Kadar FTI dalam darah
27	TBGmeasured	Apakah TBG diukur dalam darah
28	TBG	Kadar TBG dalam darah
29	Referral source	Sumber rujukan
30	binaryClass	Kelas biner

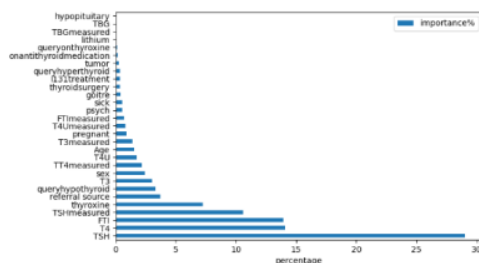
3.2. Preprocessing

Pada penelitian ini, sebelum data masuk ke tahap pemodelan, dilakukan terlebih dahulu proses pra-pemrosesan data. Pada tahapan cleaning data dilakukan proses pengisian melengkapi data yang hilang menggunakan rata-rata data. Berikut hasil data yang sudah di *cleaning* dalam dalam gambar 2.

Age	0
sex	0
thyroxine	0
queryonthyroxine	0
onantithyroidmedication	0
sick	0
pregnant	0
thyroidsurgery	0
I131treatment	0
queryhypothyroid	0
queryhyperthyroid	0
lithium	0
goitre	0
tumor	0
hypopituitary	0
psych	0
TSHmeasured	0
TSH	0
T3measured	0
T3	0
TT4measured	0
T4	0
T4Umeasured	0
T4U	0
FTImeasured	0
...	0
TBGmeasured	0
TBG	0
referral source	0
binaryClass	0

Gambar 2. Cleaning

Selanjutnya, pada tahapan selection data dilakukan untuk menunjukkan kontribusi relatif setiap fitur dalam pembuatan keputusan. Semakin tinggi nilai kepentingan suatu fitur, semakin besar pengaruh fitur tersebut dalam model. Dengan memahami nilai kepentingan masing-masing fitur dapat diidentifikasi fitur mana yang paling signifikan, sehingga dapat meningkatkan akurasi dan memiliki dampak lebih besar terhadap hasil akhir dari model. Hasil data yang sudah di selection ditampilkan dalam gambar 3.



Gambar 3. Selection

Pada tahapan transformation data, dilakukan perubahan data kategorikal menjadi numerik dengan mengonversi nilai *false* dan *true* menjadi "0" dan "1". Hasil data yang sudah di transformation dapat dilihat dalam Gambar 4.

Age	int64
sex	int32
thyroxine	int32
queryonthyroxine	int32
onantithyroidmedication	int32
sick	int32
pregnant	int32
thyroidsurgery	int32
I131treatment	int32
queryhypothyroid	int32
queryhyperthyroid	int32
lithium	int32
goitre	int32
tumor	int32
hypopituitary	int32
psych	int32
TSHmeasured	int32
TSH	float64
T3measured	int32
T3	float64
TT4measured	int32
T4	float64
T4Umeasured	int32
T4U	float64
FTImeasured	int32
...	
TBG	int32
referral source	int32
binaryClass	int32

Gambar 4. Transformation

3.3.Split Data

Setelah melalui tahapan-tahapan tersebut, dataset kemudian dibagi menjadi data pelatihan dan data pengujian. Dalam penelitian ini, pembagian dataset dilakukan dengan dua cara. Pembagian pertama adalah 80/20, di mana 80% digunakan untuk data pelatihan dan 20% untuk data pengujian. Pembagian kedua adalah 70/30, dengan 70% untuk data pelatihan dan 30% untuk data pengujian. Kedua pembagian ini bertujuan untuk menguji kinerja model dalam melakukan klasifikasi. Berikut adalah pembagian data yang dilakukan dalam tabel 2.

TABEL III. SPLIT DATA

Data Training	Data Testing
70%	30%
80%	20%

3.4.Klasifikasi Random Forest

Dalam penelitian ini, metode random forest digunakan untuk melakukan klasifikasi penyakit tiroid. Pada random forest, terdapat beberapa parameter yang diuji untuk meningkatkan kinerja dan mengoptimalkan model. Berikut adalah pengujian dengan pengaturan parameter random forest yang telah dilakukan untuk klasifikasi penyakit tiroid pada tabel 3 dan tabel 4.

TABEL IV. PENGUJIAN PERTAMA

Split Data 80/20			
Pengujian	1	2	3
5 estimator	50	100	150
max_depth	5	10	15
min_samples_leaf	5	10	15

random_state	42	42	42
Accuracy	95%	98%	99%

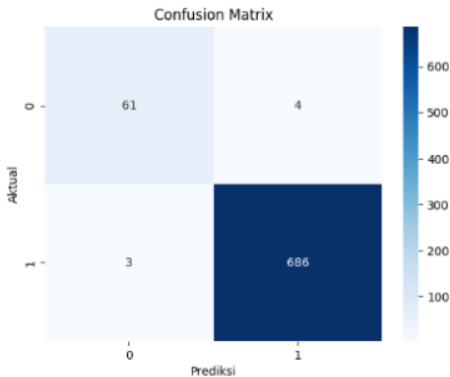
TABEL V. PENGUJIAN KEDUA

Split Data 70/30			
Pengujian	4	5	6
estimator	50	100	150
max_depth	5	10	15
min_samples_leaf	5	10	15
random_state	42	42	42
Accuracy	97%	98%	98%

Pada pengujian pertama dan kedua, diterapkan berbagai parameter yaitu ($n_estimators$) jumlah pohon keputusan diuji dengan nilai (50, 100, 150), (max_depth) kedalaman maksimum pohon keputusan diuji dengan nilai (5, 10, 15), ($min_samples_leaf$) jumlah minimum sampel diuji dengan nilai (5, 10, 15), serta $random_state$ dengan nilai 42. Setiap kombinasi parameter ini diuji untuk mengevaluasi pengaruhnya terhadap akurasi model. Hasil pengujian menunjukkan bahwa variasi parameter menghasilkan tingkat akurasi yang beragam. Hal ini menggaris bawahi pentingnya pemilihan parameter yang tepat untuk meningkatkan kinerja model. Dengan mengoptimalkan nilai $n_estimators$, max_depth , dan $min_samples_leaf$, model dapat mencapai tingkat akurasi yang lebih tinggi, menunjukkan betapa signifikan dampak tuning parameter terhadap hasil akhir.

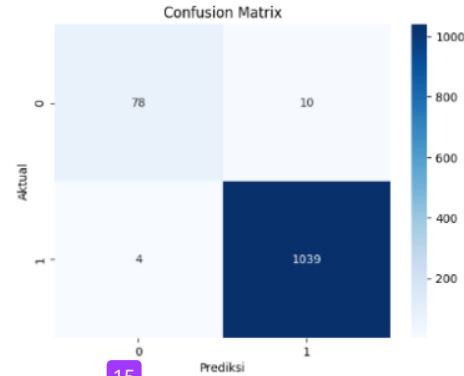
3.5. Evaluasi

Pada penelitian ini, dilakukan 6 pengujian menggunakan metode klasifikasi random forest. Pengujian ke-3 dengan pembagian data 80/20 menghasilkan akurasi terbaik. Berikut adalah *confusion matrix* untuk pengujian ke-3 yang ditampilkan dalam gambar.



Gambar 5. Confusion matrix pengujian ke-3

Untuk pembagian data 70/30, akurasi terbaik dicapai pada pengujian ke-6 dengan akurasi sebesar 98%. Pengujian kedua ini menggunakan parameter yang sama dengan pengujian pertama. Hasil dari pengujian ini lebih rendah dibandingkan pengujian pertama karena jumlah data pelatihan yang berkurang dan data pengujian yang bertambah. Berikut adalah *confusion matrix* untuk pengujian ke-6 yang ditampilkan dalam gambar 3.



Gambar 6. Confusion Matrix pengujian ke-6

4. Kesimpulan dan Saran

Berdasarkan penelitian yang telah dilakukan, dapat disimpulkan bahwa metode random forest efektif untuk klasifikasi diagnosis penyakit tiroid. Penelitian ini melibatkan 6 pengujian, dengan hasil terbaik diperoleh pada pengujian ketiga dalam pembagi data 80/20 yang mencapai akurasi 99%. Ini lebih tinggi dibandingkan dengan pembagian data 70/30 yang mencapai akurasi 98%. Perubahan parameter seperti max_depth dan $min_samples_leaf$ dapat mempengaruhi akurasi model random forest.

Saran untuk peneliti selanjutnya, agar menggunakan teknik SMOTE (Synthetic Minority Over-sampling Technique) untuk menyeimbangkan data. Teknik SMOTE ini penting karena membantu mengatasi masalah ketidakseimbangan kelas dengan menghasilkan sampel sintetis dari kelas minoritas. Dengan penerapan SMOTE, model pembelajaran mesin dapat mengenali pola dari kedua kelas dengan lebih efektif, sehingga meningkatkan akurasi dalam diagnosa penyakit tiroid.

1214 KLASIFIKASI DIAGNOSA PENYAKIT TIROID MENGUNAKAN METODE RANDOM FOREST

ORIGINALITY REPORT

18%

SIMILARITY INDEX

PRIMARY SOURCES

- | | | |
|---|---|---------------|
| 1 | e-journal.stmiklombok.ac.id
Internet | 60 words — 2% |
| 2 | Hendriyo Mokodompit, Nurnaningsih Nico Abdul, Elvie Fatmah Mokodongan. "PONDOK PESANTREN MODERN DARUL MADINAH WONOSARI KABUPATEN BOALEMO DENGAN PENDEKATAN ARSITEKTUR TROPIS", JAMBURA Journal of Architecture, 2024
Crossref | 36 words — 1% |
| 3 | eprints.binadarma.ac.id
Internet | 30 words — 1% |
| 4 | Abdul Mizwar A. Rahim, Inggrid Yanuar Risca Pratiwi, Muhammad Ainul Fikri. "Klasifikasi Penyakit Jantung Menggunakan Metode Synthetic Minority Over-Sampling Technique Dan Random Forest Clasifier", Indonesian Journal of Computer Science, 2023
Crossref | 28 words — 1% |
| 5 | github.com
Internet | 20 words — 1% |
| 6 | aaltodoc.aalto.fi
Internet | 19 words — 1% |
| 7 | eprints.ums.ac.id | |

17 words — 1%

8

repository.its.ac.id

Internet

17 words — 1%

9

text-id.123dok.com

Internet

17 words — 1%

10

Risfan Novrian, Tia Agustiani, Muhamad Fikri, Moch Fajar Hikmatulloh, Muhammad Erlangga Gunawan, Uus Firdaus. "Penerapan Algoritma Random Forest dalam Prediksi Status Penerima PIP pada Siswa: Studi Kasus pada SMK Amaliah 1", Karimah Tauhid, 2024

Crossref

15 words — 1%

11

labdas.si.fti.unand.ac.id

Internet

15 words — 1%

12

francis-press.com

Internet

14 words — 1%

13

www.researchgate.net

Internet

13 words — 1%

14

www.semanticscholar.org

Internet

13 words — 1%

15

etheses.uin-malang.ac.id

Internet

12 words — < 1%

16

pels.umsida.ac.id

Internet

11 words — < 1%

17

pdfs.semanticscholar.org

Internet

10 words — < 1%

18	repository.maranatha.edu Internet	10 words — < 1%
19	Guntur Budi Herwanto. "Document Clustering Dengan Latent Dirichlet Allocation dan Ward Hierarichal Clustering", Pseudocode, 2018 Crossref	9 words — < 1%
20	gunma-u.repo.nii.ac.jp Internet	9 words — < 1%
21	jhbmi.ir Internet	9 words — < 1%
22	publikasiilmiah.ums.ac.id Internet	9 words — < 1%
23	www.slideshare.net Internet	9 words — < 1%
24	wykoz8.tatestreetart.com Internet	9 words — < 1%
25	Asnir Dalle, Vivi Peggie Rantung, Cindy Munaiseche. "Klasifikasi Berita Hoax Pembagian Kuota Internet Menggunakan Algoritma Modified K-Nearest Neighbor", Jurnal Linguistik Komputasional (JLK), 2023 Crossref	8 words — < 1%
26	konsultasiskripsi.com Internet	8 words — < 1%
27	pubs2.ascee.org Internet	8 words — < 1%
28	www.riss.or.kr Internet	8 words — < 1%

29

ejournal.upbatam.ac.id
Internet

7 words — < 1%

30

[Ikhlusul Amalia Rahmi, Farit Mochamad Afendi, Anang Kurnia. "Metode AdaBoost dan Random Forest untuk Prediksi Peserta JKN-KIS yang Menunggak", Jambura Journal of Mathematics, 2023](#)
Crossref

6 words — < 1%

31

[Xuejing DU, Wei Zhao. "Risky lane-changing behavior recognition based on Stacking ensemble learning on snowy and icy surfaces", Research Square Platform LLC, 2024](#)
Crossref Posted Content

6 words — < 1%

EXCLUDE QUOTES OFF
EXCLUDE BIBLIOGRAPHY OFF

EXCLUDE SOURCES OFF
EXCLUDE MATCHES OFF