

Klasifikasi Penyakit Ginjal Kronis pada Data Tidak Seimbang Menggunakan K-Nearest Neighbor Berbasis Seleksi Fitur Mutual Information dan GridSearchCV

Mirza Afif Pradivta^{1,*}, Solikhun², Timbo Faritcan P. Siallagan³

¹ Information System, STIKOM Tunas Bangsa, Pematangsiantar, Indonesia

² Informatics Engineering, STIKOM Tunas Bangsa, Pematangsiantar, Indonesia

³ Faculty of Engineering, Computer and Network Engineering, Universitas Mandiri, Jawa Barat, Indonesia

Email: ^{1,*}mirza.afif27@gmail.com, ²solikhun@amiktunasbangsa.ac.id, ³timbofaritcansiallagan@gmail.com

Email Penulis Korespondensi: mirza.afif27@gmail.com

Abstrak—Penyakit Ginjal Kronis atau Chronic Kidney Disease (CKD) merupakan penyakit progresif yang menurunkan fungsi ginjal secara bertahap dan memerlukan deteksi dini untuk mengurangi risiko komplikasi berat. Machine Learning dapat mendukung klasifikasi CKD berdasarkan atribut klinis, tetapi dataset medis sering memiliki missing value, gabungan fitur numerik dan kategorikal, serta distribusi kelas yang tidak seimbang. Penelitian ini bertujuan untuk mengoptimasi algoritma K-Nearest Neighbor (KNN) pada klasifikasi CKD menggunakan seleksi fitur Mutual Information dan GridSearchCV. Dataset terdiri atas 400 sampel dengan dua kelas, yaitu 250 sampel CKD dan 150 sampel not CKD. Tahapan penelitian meliputi pembersihan data, penanganan missing value, encoding fitur kategorikal, normalisasi fitur numerik menggunakan MinMaxScaler, seleksi fitur menggunakan SelectKBest dengan Mutual Information, serta optimasi hyperparameter menggunakan GridSearchCV. Evaluasi dilakukan melalui hold-out testing dan 10-Fold Cross Validation. Hasil evaluasi hold-out menunjukkan bahwa model KNN dengan GridSearchCV mencapai accuracy, precision, recall, F1-score, dan AUC sebesar 100,00% pada data uji. Untuk memastikan bahwa hasil tersebut tidak hanya bergantung pada satu pembagian data, evaluasi juga dilakukan menggunakan 10-Fold Cross Validation. Evaluasi menggunakan 10-Fold Cross Validation menghasilkan rata-rata akurasi sebesar 99,25% pada model KNN dengan GridSearchCV, yang menunjukkan bahwa performa model tetap konsisten pada berbagai pembagian data. Sementara itu, skenario KNN + Feature Selection + GridSearchCV memperoleh accuracy 98,75%, recall 100,00%, dan F1-score 99,01%, sehingga tetap menunjukkan performa kompetitif dengan subset fitur yang lebih ringkas. Hasil penelitian menunjukkan bahwa penerapan GridSearchCV meningkatkan performa KNN pada dataset yang digunakan, sedangkan Mutual Information membantu memilih fitur-fitur yang relevan sehingga model tetap memberikan performa yang baik dengan jumlah fitur yang lebih ringkas.

Kata Kunci: Penyakit Ginjal Kronis, klasifikasi, GridSearchCV, KNN, Mutual Information.

Abstract—Chronic Kidney Disease (CKD) is a progressive disease characterized by a gradual decline in kidney function and requires early detection to reduce the risk of severe complications. Machine learning has been widely applied to support CKD classification based on clinical attributes; however, medical datasets often contain missing values, a combination of numerical and categorical features, and class imbalance. This study aims to evaluate the performance of the K-Nearest Neighbor (KNN) algorithm for CKD classification using Mutual Information feature selection and GridSearchCV. The dataset consisted of 400 samples, including 250 CKD cases and 150 non-CKD cases. The proposed methodology included data cleaning, missing value imputation, categorical feature encoding, numerical feature normalization using MinMaxScaler, feature selection using SelectKBest with Mutual Information, and hyperparameter tuning using GridSearchCV. Model performance was evaluated using hold-out testing and 10-fold cross-validation. The hold-out evaluation showed that the KNN model with GridSearchCV achieved 100.00% accuracy, precision, recall, F1-score, and AUC on the test set. To ensure that this result was not dependent on a single train-test split, additional evaluation was conducted using 10-fold cross-validation. The cross-validation results yielded an average accuracy of 99.25% for the KNN model with GridSearchCV, indicating consistent performance across different data partitions. Meanwhile, the KNN model with Mutual Information feature selection and GridSearchCV achieved 98.75% accuracy, 100.00% recall, and a 99.01% F1-score, demonstrating competitive performance while using a more compact feature subset. The findings indicate that the application of GridSearchCV improved the performance of the KNN model on the dataset used, while Mutual Information contributed to selecting relevant features, enabling the model to maintain strong classification performance with a reduced number of features.

Keywords: Chronic Kidney Disease; Classification; GridSearchCV; K-Nearest Neighbor; Mutual Information

1. PENDAHULUAN

Penyakit Ginjal Kronis atau Chronic Kidney Disease (CKD) merupakan penyakit progresif yang ditandai dengan penurunan fungsi ginjal secara bertahap dalam periode tertentu. Ginjal memiliki peran penting dalam menyaring limbah metabolik, menjaga keseimbangan cairan, mengatur elektrolit, serta mendukung pengeluaran zat sisa melalui urine. Ketika fungsi ginjal terganggu, tubuh tidak dapat mempertahankan keseimbangan metabolisme secara optimal. Kondisi ini dapat menyebabkan penumpukan zat berbahaya dan meningkatkan risiko komplikasi serius, terutama apabila penyakit tidak terdeteksi pada tahap awal[1]-[4].

Deteksi dini CKD menjadi penting karena gejalanya sering berkembang secara perlahan dan tidak selalu terlihat pada tahap awal. Banyak pasien baru menyadari penyakit tersebut ketika fungsi ginjal telah mengalami penurunan yang cukup berat. Identifikasi lebih awal dapat membantu memperlambat perkembangan penyakit, mengurangi risiko komplikasi, serta mendukung pengambilan keputusan medis yang lebih tepat [5]–[8]. Oleh karena itu, diperlukan pendekatan komputasional untuk membantu identifikasi status CKD berdasarkan atribut klinis secara lebih cepat dan konsisten. Perkembangan Machine Learning memberikan peluang untuk mendukung klasifikasi

penyakit berdasarkan data medis. Machine Learning dapat mempelajari pola dari atribut pasien dan menghasilkan model yang mampu membedakan kondisi pasien berdasarkan karakteristik klinis. Pada klasifikasi CKD, metode supervised learning dapat digunakan untuk mengklasifikasikan pasien ke dalam kelas CKD dan not CKD menggunakan atribut seperti tekanan darah, hemoglobin, serum creatinine, kadar glukosa darah, serta indikator klinis lainnya [9]–[16].

Salah satu algoritma Machine Learning yang dapat digunakan untuk klasifikasi adalah K-Nearest Neighbor (KNN). KNN menentukan kelas suatu sampel berdasarkan sampel berlabel terdekat dalam ruang fitur. Algoritma ini sederhana, mudah diterapkan[17], dan dapat digunakan pada berbagai permasalahan klasifikasi[16]. Pada klasifikasi CKD, KNN dapat memanfaatkan kemiripan antar atribut klinis pasien untuk menentukan apakah pasien termasuk dalam kelas CKD atau not CKD [18]–[21].

Meskipun KNN memiliki konsep yang sederhana, algoritma ini memiliki beberapa keterbatasan. KNN sensitif terhadap skala fitur karena proses klasifikasinya bergantung pada perhitungan jarak. Fitur dengan rentang nilai yang lebih besar dapat mendominasi perhitungan jarak apabila normalisasi tidak diterapkan dengan baik. Selain itu, atribut yang kurang relevan dapat menurunkan kualitas prediksi karena mengganggu pengukuran kemiripan antar sampel. Hyperparameter seperti jumlah tetangga terdekat, strategi pembobotan, dan metrik jarak juga memengaruhi performa klasifikasi[15]- [21]-[23].

Beberapa penelitian sebelumnya telah menerapkan KNN[20] dan algoritma Machine Learning lain untuk klasifikasi CKD. Penelitian tersebut menunjukkan bahwa KNN dapat mengklasifikasikan kasus CKD, tetapi performanya masih dapat ditingkatkan melalui optimasi, preprocessing yang tepat, dan evaluasi yang cermat [5], [11], [12]. Penelitian lain juga menunjukkan bahwa GridSearchCV dapat meningkatkan pemilihan parameter pada model Machine Learning dan tugas klasifikasi berbasis KNN [10]-[14]. Mutual Information dipilih sebagai metode seleksi fitur karena mampu mengukur tingkat ketergantungan antara setiap fitur dan kelas target tanpa mengasumsikan hubungan linier maupun distribusi data tertentu. Karakteristik tersebut sesuai dengan dataset CKD yang terdiri atas kombinasi atribut numerik dan kategorikal serta mengandung missing value. Dibandingkan beberapa metode seleksi fitur lain, Mutual Information mampu mengidentifikasi fitur yang memiliki kontribusi informasi tinggi terhadap proses klasifikasi sehingga relevan untuk mendukung kinerja algoritma K-Nearest Neighbor yang sangat bergantung pada kualitas fitur masukan.

Berdasarkan penelitian sebelumnya, masih terdapat celah penelitian pada integrasi seleksi fitur dan hyperparameter tuning untuk klasifikasi CKD. Penelitian terdahulu umumnya berfokus pada perbandingan model secara langsung atau optimasi parameter, sedangkan peran seleksi fitur Mutual Information dalam mengurangi atribut yang kurang informatif belum ditekankan secara memadai. Pada dataset medis yang memiliki fitur numerik, fitur kategorikal, missing value, dan distribusi kelas tidak seimbang ringan hingga sedang, penggunaan seluruh atribut tanpa seleksi fitur dapat menyebabkan model mempelajari informasi yang kurang relevan. Kondisi tersebut menjadi penting karena KNN sangat bergantung pada perhitungan jarak, sehingga keberadaan atribut yang kurang informatif dapat memengaruhi kualitas klasifikasi.

Penelitian ini bertujuan untuk mengoptimasi algoritma KNN dalam klasifikasi CKD pada data tidak seimbang menggunakan seleksi fitur Mutual Information dan GridSearchCV. Kontribusi penelitian ini mencakup preprocessing data medis, pemilihan fitur relevan menggunakan Mutual Information, optimasi hyperparameter KNN menggunakan GridSearchCV, serta perbandingan empat skenario model, yaitu KNN Default, KNN dengan Feature Selection, KNN dengan GridSearchCV, dan KNN dengan Feature Selection dan GridSearchCV. Struktur artikel terdiri atas Introduction, Method, Result, Discussions, Conclusion, Conflict of Interest, Acknowledgement, dan References.

2. METODOLOGI PENELITIAN

Bagian metode menjelaskan tahapan penelitian yang digunakan untuk mengolah dataset CKD, membangun model klasifikasi, serta mengevaluasi performa model. Bagian ini berfokus pada prosedur dan metode penelitian, sedangkan hasil numerik dan interpretasinya disajikan pada bagian HASIL dan PEMBAHASAN.



Gambar 1. Diagram Alur Penelitian

Gambar 1 mengilustrasikan alur penelitian yang digunakan dalam studi ini. Proses dimulai dari pemuatan dataset CKD, preprocessing, seleksi fitur, dan pembagian data, kemudian dilanjutkan dengan empat skenario model dan evaluasi menggunakan beberapa metrik performa.

2.1 Dataset

Dataset yang digunakan dalam penelitian ini adalah dataset Chronic Kidney Disease (CKD) dari UCI Machine Learning Repository [26]. Dataset tersebut berisi data klinis pasien dengan target klasifikasi yang merepresentasikan status CKD. Dataset terdiri atas 400 record dengan 26 kolom awal, yang mencakup 25 kolom atribut dan satu kolom target. Kolom target yang digunakan dalam penelitian ini adalah classification, yang terdiri atas dua kelas, yaitu ckd untuk pasien yang terindikasi penyakit ginjal kronis dan notckd untuk pasien yang tidak terindikasi penyakit ginjal kronis.

Sebelum pemodelan dilakukan, kolom id dihapus karena hanya berfungsi sebagai atribut identitas dan tidak memberikan informasi prediktif untuk klasifikasi CKD. Dengan demikian, 24 fitur digunakan dalam proses klasifikasi, sedangkan kolom classification digunakan sebagai target. Dataset memiliki fitur numerik dan kategorikal sehingga diperlukan preprocessing sebelum data dapat diproses oleh model Machine Learning.

Tabel 1. Sampel Data CKD

No.	Age	BP	SG
1	48.0	80.0	1.020
2	7.0	50.0	1.020
3	62.0	80.0	1.010
4	48.0	70.0	1.005
5	51.0	80.0	1.010
6	60.0	90.0	1.015
7	68.0	70.0	1.010
8	24.0	–	1.015
9	52.0	100.0	1.015
10	53.0	90.0	1.020
...	...	...	...
400	58.0	80.0	1.025

Tabel 1 menampilkan sampel data CKD yang digunakan dalam penelitian ini. Sampel data tersebut memperlihatkan sebagian atribut numerik, yaitu age, bp, dan sg, sebagai gambaran struktur dataset. Dataset lengkap terdiri atas 400 record dengan 24 fitur yang digunakan dalam proses klasifikasi setelah kolom id dihapus.

Tabel 2. Distribusi Kelas Awal

Kelas	Jumlah Sampel	Persentase	Keterangan
CKD	250	62.50%	Kelas mayoritas
Not CKD	150	37.50%	Kelas minoritas
Total	400	100.00%	-

Tabel 2 menunjukkan bahwa dataset memiliki distribusi kelas yang tidak seimbang. Kelas CKD terdiri atas 250 sampel atau 62,50%, sedangkan kelas not CKD terdiri atas 150 sampel atau 37,50%. Distribusi kelas ini menunjukkan bahwa dataset memiliki ketidakseimbangan ringan hingga sedang, sehingga evaluasi model tidak boleh hanya bergantung pada akurasi.

2.2 Data Preprocessing

Data preprocessing dilakukan untuk memastikan dataset dapat diproses dengan baik oleh algoritma KNN. Tahap preprocessing meliputi pembersihan nama kolom dan nilai teks, penggantian nilai tidak valid seperti tanda tanya dan string kosong menjadi missing value, konversi kolom numerik yang terbaca sebagai teks menjadi format numerik, serta encoding kelas target. Kelas ckd dikodekan sebagai 1, sedangkan kelas notckd dikodekan sebagai 0.

Missing value pada atribut numerik ditangani menggunakan median, sedangkan missing value pada atribut kategorikal ditangani menggunakan nilai yang paling sering muncul. Fitur numerik dinormalisasi menggunakan MinMaxScaler karena KNN sensitif terhadap skala fitur. Fitur kategorikal ditransformasikan menggunakan OrdinalEncoder. Untuk mengurangi risiko data leakage, komponen preprocessing di-fit pada data latih dan kemudian diterapkan pada data uji melalui pipeline pemodelan. Dataset dibagi menjadi 80% data latih dan 20% data uji menggunakan stratified split untuk evaluasi hold-out. Selain itu, evaluasi terhadap seluruh 400 data dilakukan menggunakan 10-Fold Cross Validation dengan pendekatan stratified sehingga proporsi kelas CKD dan not CKD tetap terjaga pada setiap fold, serta setiap data menjadi bagian data uji secara bergantian.

Untuk menghindari data leakage, pembagian data dilakukan terlebih dahulu menggunakan stratified split. Selanjutnya, proses imputasi missing value, encoding fitur kategorikal, normalisasi menggunakan MinMaxScaler, dan seleksi fitur dilakukan pada data latih (fit). Parameter yang diperoleh kemudian diterapkan pada data uji (transform)

sehingga informasi dari data uji tidak digunakan selama proses pelatihan model.

2.3 Feature Selection

Feature selection dilakukan menggunakan SelectKBest dengan Mutual Information. Mutual Information mengukur ketergantungan antara setiap fitur dan kelas target serta umum digunakan untuk mengidentifikasi atribut informatif pada tugas klasifikasi. Fitur dengan skor Mutual Information yang lebih tinggi dianggap lebih relevan untuk membedakan CKD dan not CKD. Tahap ini bertujuan untuk mengurangi fitur yang tidak relevan dan membantu algoritma KNN menghitung jarak berdasarkan atribut yang lebih informatif. Mutual Information digunakan juga untuk mengukur ketergantungan antara setiap fitur input dan kelas target. Skor Mutual Information yang lebih tinggi menunjukkan bahwa suatu fitur mengandung lebih banyak informasi relevan untuk membedakan kelas CKD dan non-CKD. Dalam penelitian ini, SelectKBest diterapkan untuk memilih fitur yang paling informatif sebelum proses klasifikasi KNN. Langkah ini penting karena KNN bergantung pada perhitungan jarak, dan fitur yang tidak relevan dapat mengurangi kinerja klasifikasi dengan memperkenalkan noise yang tidak perlu ke dalam pengukuran jarak.

2.4 K-Nearest Neighbor dan GridSearchCV

K-Nearest Neighbor digunakan sebagai algoritma klasifikasi utama. Model menentukan kelas sampel baru berdasarkan kelas mayoritas dari tetangga terdekatnya. KNN mengklasifikasikan sampel pengujian dengan menghitung jaraknya ke sampel pelatihan dan menetapkan label kelas berdasarkan kelas mayoritas di antara tetangga terdekat. Dalam penelitian ini, pengukuran jarak merupakan komponen penting karena dataset CKD mengandung atribut numerik dengan rentang nilai yang berbeda. Oleh karena itu, normalisasi numerik diterapkan sebelum pelatihan model. Salah satu pengukuran jarak yang umum digunakan dalam KNN adalah jarak Euclidean, yang dapat dinyatakan sebagai berikut:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (1)$$

di mana ($d(x,y)$) merupakan jarak antara data (x) dan data (y), (x_i) merupakan nilai fitur ke- (i) pada data (x), (y_i) merupakan nilai fitur ke- (i) pada data (y), dan (n) merupakan jumlah fitur yang digunakan dalam proses klasifikasi. Semakin kecil nilai jarak yang dihasilkan, semakin tinggi tingkat kemiripan antara data uji dan data latih. Dalam penelitian ini, jarak Euclidean digunakan untuk mengukur kedekatan antar data pasien berdasarkan atribut klinis yang telah melalui proses preprocessing dan normalisasi.

Performa KNN bergantung pada beberapa hyperparameter, termasuk jumlah tetangga, metode pembobotan, dan metrik jarak. Oleh karena itu, GridSearchCV digunakan untuk mencari kombinasi hyperparameter KNN terbaik secara sistematis dalam proses pelatihan dan cross-validation.

Tabel 3. Skenario Evaluasi

Skenario	Model	Keterangan
1	KNN Default	Baseline tanpa feature selection dan tuning
2	KNN + Feature Selection	KNN dengan feature selection Mutual Information
3	KNN + GridSearchCV	Model terbaik pada hold-out testing
4	IN + Feature Selection + GridSearchCV	Skenario integrasi feature selection dan tuning

Tabel 4. Konfigurasi Parameter GridSearchCV

Parameter	Nilai yang Diuji
feature_selection__k	5, 7, 9, 11, 13, 15, 17, all
n_neighbors	1, 3, 5, ..., 29
weights	uniform, distance
metric	euclidean, manhattan, minkowski
p	1, 2, 3
Cross Validation	10-Fold Stratified K-Fold

2.5 Metrik Evaluasi

Model dievaluasi menggunakan accuracy, precision, recall, F1-score, AUC, confusion matrix, 10-Fold Cross Validation, dan Wilcoxon Signed-Rank Test. Evaluasi dilakukan melalui dua pendekatan, yaitu hold-out testing untuk menguji performa model pada data uji terpisah dan 10-Fold Cross Validation untuk memastikan seluruh 400 data dievaluasi secara bergantian sebagai data uji. Karena dataset memiliki distribusi kelas tidak seimbang ringan hingga sedang, interpretasi juga mempertimbangkan recall per kelas, specificity, balanced accuracy, dan performa tingkat makro jika diperlukan. Accuracy mengukur proporsi prediksi yang benar, precision mengukur ketepatan prediksi positif, recall mengukur kemampuan model dalam mengidentifikasi sampel positif aktual, dan F1-score

menggabungkan precision dan recall. AUC mengukur kemampuan model dalam membedakan kelas CKD dan not CKD.

$$Acc = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

$$Precision = \frac{TP}{TP + FP} \quad (3)$$

$$Recall = \frac{TP}{TP + FN} \quad (4)$$

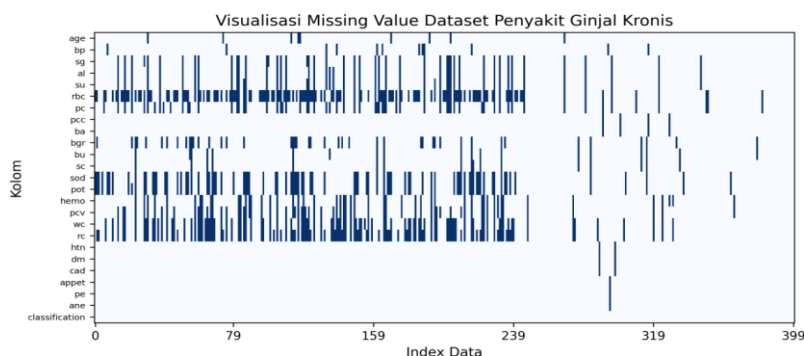
$$F1-score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (5)$$

TP menunjukkan sampel CKD yang diklasifikasikan dengan benar sebagai CKD, TN menunjukkan sampel not CKD yang diklasifikasikan dengan benar sebagai not CKD, FP menunjukkan sampel not CKD yang salah diklasifikasikan sebagai CKD, dan FN menunjukkan sampel CKD yang salah diklasifikasikan sebagai not CKD. Specificity diinterpretasikan sebagai $TN / (TN + FP)$, sedangkan balanced accuracy diinterpretasikan sebagai rata-rata dari sensitivity dan specificity.

3. HASIL DAN PEMBAHASAN

3.1 Hasil Preprocessing

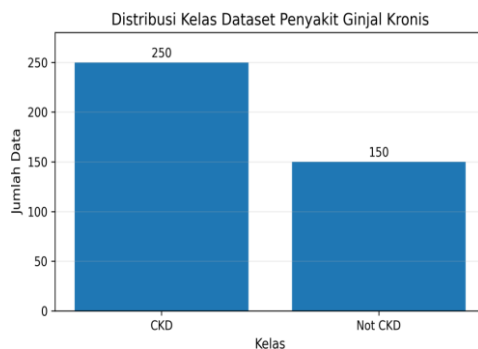
Tahap preprocessing mengidentifikasi adanya missing value pada beberapa atribut. Missing value tersebut ditangani menggunakan imputasi median untuk atribut numerik dan imputasi nilai yang paling sering muncul untuk atribut kategorikal. Visualisasi missing value ditampilkan pada Gambar 2.



Gambar 2. Plot Missing Value

Visualisasi menunjukkan bahwa missing value tersebar pada beberapa atribut klinis sehingga diperlukan proses imputasi sebelum pembangunan model. Tahapan ini bertujuan menjaga kualitas data tanpa mengurangi jumlah sampel yang digunakan.

Gambar 2 menunjukkan bahwa missing value tersebar pada beberapa atribut klinis, seperti (sebutkan atribut yang memang banyak kosong, misalnya rbc, pc, su, atau atribut lain sesuai gambar). Kondisi ini menunjukkan bahwa dataset memerlukan tahap imputasi sebelum proses klasifikasi agar seluruh sampel dapat dimanfaatkan tanpa mengurangi jumlah data. Oleh karena itu, penelitian ini menerapkan imputasi median untuk atribut numerik dan modus untuk atribut kategorikal sebagai bagian dari tahap preprocessing.

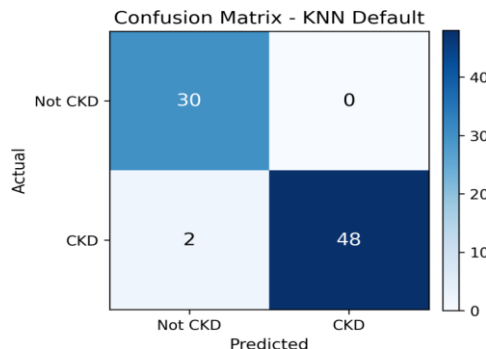


Gambar 3. Distribusi Kelas Dataset CKD

Gambar 3 menunjukkan bahwa kelas CKD lebih dominan dibandingkan kelas not CKD. Ketidakseimbangan ini mendukung penggunaan metrik evaluasi tambahan seperti precision, recall, F1-score, dan AUC.

3.2 Hasil KNN Default

Setelah preprocessing, data diuji menggunakan KNN Default sebagai model baseline. Model default digunakan untuk mengidentifikasi performa awal sebelum penerapan feature selection dan hyperparameter tuning.

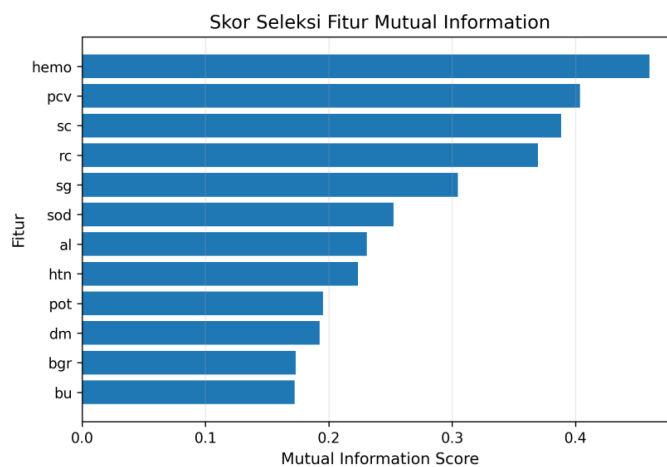


Gambar 4. Confusion Matrix KNN Default

Gambar 4 menunjukkan bahwa KNN Default berhasil mengklasifikasikan seluruh sampel not CKD dengan benar dan 48 dari 50 sampel CKD pada data uji hold-out. Dua sampel CKD salah diklasifikasikan sebagai not CKD. Hasil ini menunjukkan bahwa KNN Default sudah menghasilkan performa yang tinggi, tetapi peningkatan masih memungkinkan.

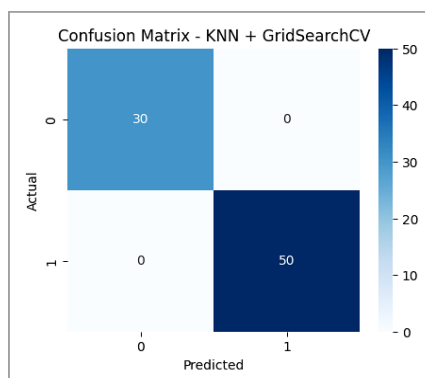
3.3 Hasil KNN dengan Feature Selection

Skenario kedua menerapkan seleksi fitur Mutual Information sebelum proses klasifikasi KNN. Fitur yang terpilih adalah sg, al, sc, sod, hemo, pcv, rc, htn, dan dm. Penggunaan fitur-fitur tersebut meningkatkan akurasi pengujian dari 97,50% menjadi 98,75%.



Gambar 5. Skor Fitur Mutual Information

3.4 Hasil KNN dengan GridSearchCV

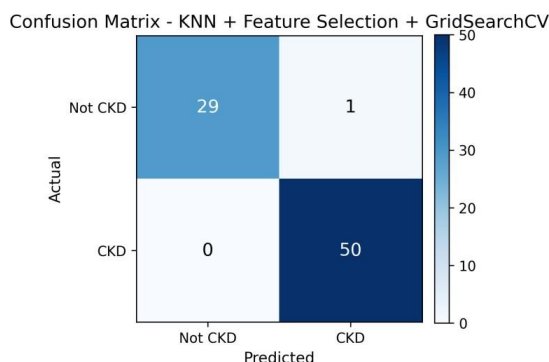


Gambar 6. Confusion Matrix KNN dengan GridSearchCV

Skenario ketiga menerapkan GridSearchCV untuk mencari kombinasi hyperparameter terbaik pada KNN. Berdasarkan hasil pengujian hold-out, skenario KNN + GridSearchCV memperoleh performa tertinggi dengan accuracy, precision, recall, F1-score, dan AUC sebesar 100,00%. Confusion matrix menunjukkan bahwa seluruh sampel pada data uji hold-out berhasil diklasifikasikan dengan benar[19]. Dengan demikian, KNN + GridSearchCV ditetapkan sebagai model dengan performa terbaik pada pengujian data uji. Namun, karena performa sempurna pada data uji dapat dipengaruhi oleh ukuran dataset dan komposisi pembagian data, hasil ini tetap diinterpretasikan bersama 10-Fold Cross Validation yang mengevaluasi seluruh 400 data secara bergantian.

3.5 Hasil KNN dengan Feature Selection dan GridSearchCV

Skenario keempat menggabungkan seleksi fitur Mutual Information dan GridSearchCV. Skenario ini digunakan untuk mengevaluasi pengaruh integrasi seleksi fitur dan optimasi hyperparameter terhadap performa KNN. Konfigurasi parameter terbaik adalah *feature_selection* $k = 9$, *metric* = manhattan, *n_neighbors* = 13, dan *weights* = uniform. Meskipun skenario ini menghasilkan performa yang kompetitif dengan accuracy 98,75%, recall 100,00%, dan F1-score 99,01%, hasilnya masih berada di bawah skenario KNN + GridSearchCV yang memperoleh accuracy 100,00% pada pengujian hold-out. Oleh karena itu, KNN + GridSearchCV ditetapkan sebagai model dengan performa terbaik, sedangkan skenario KNN + Feature Selection + GridSearchCV menunjukkan bahwa seleksi fitur dapat menghasilkan performa tinggi dengan subset fitur yang lebih ringkas.

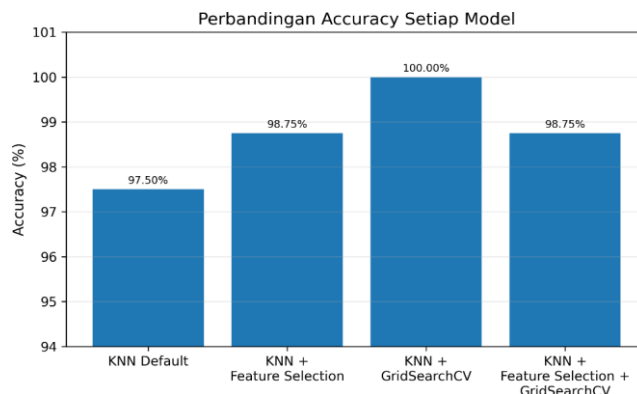


Gambar 7. Confusion Matrix KNN dengan Feature Selection dan GridSearchCV

Tabel 5. Perbandingan Performa Model pada Data Uji Hold-Out

Model	Akurasi	Presisi	Recall	F1-score	AUC
KNN Default	97.50%	100.00%	96.00%	97.96%	100.00%
KNN + Feature Selection	98.75%	100.00%	98.00%	98.99%	99.97%
KNN + GridSearchCV	100.00%	100.00%	100.00%	100.00%	100.00%
KNN + Feature Selection + GridSearchCV	98.75%	98.04%	100.00%	99.01%	99.97%

Tabel 5 menunjukkan bahwa skenario KNN + GridSearchCV memperoleh performa tertinggi pada data uji hold-out dengan accuracy, precision, recall, F1-score, dan AUC sebesar 100,00%. Hasil ini menunjukkan bahwa optimasi hyperparameter melalui GridSearchCV mampu meningkatkan performa KNN dibandingkan model default dan skenario feature selection. Sementara itu, KNN + Feature Selection + GridSearchCV tetap menghasilkan performa kompetitif dengan accuracy 98,75% dan recall 100,00%, tetapi tidak melampaui performa KNN + GridSearchCV pada data uji hold-out.



Gambar 8. Perbandingan Akurasi Setiap Model

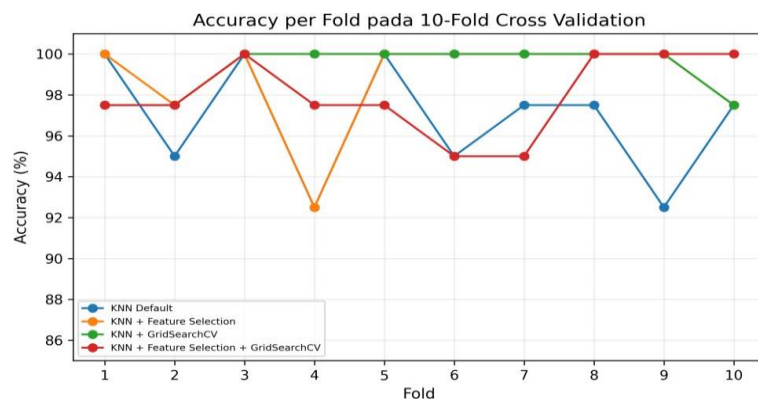
Meskipun model KNN dengan GridSearchCV memperoleh akurasi 100,00% pada skenario hold-out testing, hasil tersebut tidak serta-merta menunjukkan kemampuan generalisasi yang sempurna. Nilai tersebut dapat dipengaruhi oleh karakteristik dataset maupun pembagian data pada skenario hold-out. Oleh karena itu, dilakukan evaluasi menggunakan 10-Fold Cross Validation. Hasil rata-rata akurasi sebesar 99,25% menunjukkan bahwa performa model tetap konsisten pada berbagai pembagian data sehingga hasil yang diperoleh tidak hanya bergantung pada satu pembagian data.

3.6 Hasil 10-Fold Cross Validation

Untuk mengevaluasi kestabilan model dan memenuhi pengujian terhadap seluruh data, 10-Fold Cross Validation diterapkan. Pada skema ini, seluruh 400 data digunakan secara bergantian sebagai data uji pada setiap fold, sehingga evaluasi tidak hanya bergantung pada satu pembagian train-test split. Nilai akurasi secara rinci pada setiap fold ditampilkan pada Tabel 6, sedangkan visualisasi akurasi per fold ditampilkan pada Gambar 9. Berdasarkan hasil tersebut, KNN + GridSearchCV memperoleh rata-rata akurasi tertinggi sebesar 99,25% dengan standar deviasi 1,21%.

Tabel 6. Akurasi per Fold pada 10-Fold Cross Validation terhadap 400 Data

Fold	KNN Default	KNN + FS	KNN + GridCV	KNN + FS + GridCV
1	100.0%	100.0%	97.5%	97.5%
2	95.0%	97.5%	97.5%	97.5%
3	100.0%	100.0%	100.0%	100.0%
4	92.5%	92.5%	100.0%	97.5%
5	100.0%	100.0%	100.0%	97.5%
6	95.0%	100.0%	100.0%	95.0%
7	97.5%	100.0%	100.0%	95.0%
8	97.5%	100.0%	100.0%	100.0%
9	92.5%	100.0%	100.0%	100.0%
10	97.5%	97.5%	97.5%	100.0%
Rata-rata	96.75%	98.75%	99.25%	98.00%
Std Akurasi	2.90%	2.43%	1.21%	1.97%



Gambar 9. Akurasi per Fold pada 10-Fold Cross Validation terhadap 400 Data

3.7 Hasil Wilcoxon Signed-Rank Test

Wilcoxon Signed-Rank Test digunakan sebagai evaluasi statistik tambahan terhadap perbedaan performa model berbasis fold. Pengujian ini digunakan untuk melihat apakah peningkatan performa yang diamati pada skenario optimasi memiliki perbedaan yang signifikan secara statistik dibandingkan KNN Default. Nilai p-value di bawah 0,05 menunjukkan perbedaan yang signifikan secara statistik, sedangkan p-value yang sama dengan atau di atas 0,05 menunjukkan bahwa peningkatan yang diamati belum signifikan secara statistik pada tingkat 0,05.

Tabel 7. Hasil Wilcoxon Signed-Rank Test

Metrik	P-value	Keputusan
Akurasi	0.19531	Tidak signifikan karena $p \geq 0,05$
F1-score	0.05078	Tidak signifikan karena $p \geq 0,05$

3.8 Pembahasan

Bagian pembahasan menginterpretasikan hasil yang diperoleh pada bagian sebelumnya. Bagian ini berfokus pada penjelasan peningkatan performa, peran seleksi fitur Mutual Information, interpretasi ROC curve, perbandingan dengan penelitian rujukan, dan keterbatasan penelitian.

3.8.1 Analisis Performa Model

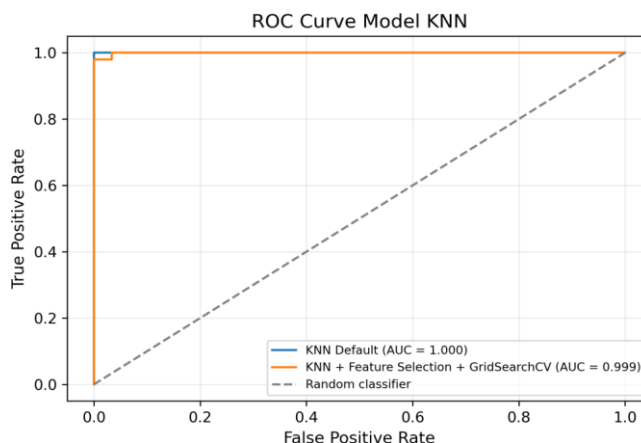
Perbandingan performa model menunjukkan bahwa KNN Default sudah menghasilkan akurasi yang tinggi pada dataset CKD. Kondisi ini menunjukkan bahwa dataset memiliki pola yang relatif dapat dibedakan antara kelas CKD dan not CKD. Berdasarkan pengujian hold-out, KNN + GridSearchCV menjadi model dengan performa terbaik karena memperoleh accuracy, precision, recall, F1-score, dan AUC sebesar 100,00%. Hasil ini menunjukkan bahwa optimasi hyperparameter memiliki pengaruh paling kuat terhadap peningkatan performa KNN pada dataset CKD.

Walaupun judul penelitian menekankan penggunaan seleksi fitur Mutual Information dan GridSearchCV, hasil eksperimen menunjukkan bahwa performa tertinggi diperoleh oleh skenario KNN + GridSearchCV. Sementara itu, skenario KNN + Feature Selection + GridSearchCV tetap dipertahankan sebagai bagian penting dari penelitian karena menunjukkan pengaruh Mutual Information dalam memilih subset fitur yang lebih ringkas dan relevan secara klinis. Skenario tersebut memperoleh accuracy 98,75%, recall 100,00%, dan F1-score 99,01% pada hold-out testing. Dengan demikian, GridSearchCV berperan paling dominan dalam peningkatan akurasi, sedangkan Mutual Information berperan dalam penyederhanaan fitur dan peningkatan interpretabilitas model. Meskipun KNN + GridSearchCV memperoleh akurasi 100,00% pada hold-out testing, hasil tersebut tetap perlu ditafsirkan bersama 10-Fold Cross Validation karena data uji hold-out hanya mewakili satu komposisi pembagian data. Berdasarkan hasil cross-validation, KNN + GridSearchCV tetap memperoleh rata-rata akurasi tertinggi sebesar 99,25% dengan standar deviasi 1,21%, sehingga hasil utama tidak hanya bergantung pada satu pembagian data uji.

3.8.2 Analisis Fitur Mutual Information

Fitur yang terpilih meliputi sg, al, sc, sod, hemo, pcv, rc, htn, dan dm. Atribut-atribut tersebut relevan secara klinis terhadap klasifikasi CKD. Serum creatinine dan specific gravity berkaitan erat dengan fungsi ginjal, sedangkan hemoglobin, packed cell volume, dan red blood cell count berkaitan dengan kondisi darah yang dapat menyertai CKD. Hypertension dan diabetes mellitus juga merupakan faktor klinis penting karena sering menjadi komorbiditas yang berkaitan dengan progresivitas CKD [2], [7].

3.8.3 Analisis ROC Curve



Gambar 10. ROC Curve Model KNN

Gambar 10 menunjukkan ROC curve dari model yang diuji. Nilai AUC mendekati 1,00, yang menunjukkan bahwa model memiliki kemampuan kuat dalam membedakan kelas CKD dan not CKD. Pada pengujian hold-out, KNN + GridSearchCV memperoleh AUC sebesar 100,00%, sedangkan KNN + Feature Selection + GridSearchCV memperoleh AUC sebesar 99,97%. Meskipun demikian, interpretasi ROC perlu dipertimbangkan bersama hasil cross-validation karena ukuran dataset yang terbatas.

3.8.4 Perbandingan dengan Penelitian Rujukan

Dibandingkan dengan penelitian sebelumnya yang menerapkan KNN dan optimasi parameter untuk klasifikasi penyakit, penelitian ini mengevaluasi GridSearchCV sebagai optimasi utama dan menambahkan seleksi fitur Mutual Information sebagai komponen analisis tambahan [10], [14]. GridSearchCV memberikan performa terbaik pada pengujian hold-out, sedangkan Mutual Information membantu mengidentifikasi fitur yang relevan secara klinis dan mengurangi pengaruh atribut yang kurang informatif. Perbandingan kontribusi penelitian ditampilkan pada Tabel 8.

Tabel 8. Perbandingan Kontribusi Penelitian

Aspek	Penelitian Rujukan	Penelitian Ini	Kontribusi
Algoritma	KNN + GridSearchCV	KNN + GridSearchCV dan KNN + Mutual Information + GridSearchCV	Menggabungkan tuning parameter dan feature selection pada evaluasi KNN
Feature Selection	Tidak dijelaskan	Mutual Information	Memilih fitur relevan untuk

Aspek	Penelitian Rujukan	Penelitian Ini	Kontribusi
	secara khusus		klasifikasi CKD
Tuning	GridSearchCV	GridSearchCV	Menggunakan pencarian hyperparameter secara sistematis
Evaluasi	Accuracy, F1-score, AUC	Accuracy, Precision, Recall, F1-score, AUC, 10-Fold CV, Wilcoxon, ROC Curve	Memberikan evaluasi yang lebih lengkap
Karakteristik Dataset	Dataset CKD	Dataset CKD dengan 250 sampel CKD dan 150 sampel not CKD	Mengevaluasi model pada data tidak seimbang ringan hingga sedang

3.8.5 Keterbatasan Penelitian

Meskipun model terbaik menghasilkan performa yang tinggi, penelitian ini masih memiliki keterbatasan. Pertama, dataset hanya terdiri atas 400 sampel dari repositori publik sehingga model perlu divalidasi lebih lanjut menggunakan dataset klinis yang lebih besar dan beragam. Kedua, hasil akurasi 100,00% pada hold-out testing perlu diinterpretasikan bersama 10-Fold Cross Validation karena hold-out testing hanya menggunakan satu komposisi pembagian data. Ketiga, Wilcoxon Signed-Rank Test menunjukkan bahwa peningkatan belum signifikan secara statistik pada tingkat signifikansi 0,05 karena performa baseline KNN sudah sangat tinggi. Keempat, dataset memiliki missing value sehingga performa dapat dipengaruhi oleh strategi imputasi. Kelima, distribusi kelas berada pada tingkat tidak seimbang ringan hingga sedang sehingga penelitian selanjutnya perlu mengevaluasi metrik tambahan yang peka terhadap ketidakseimbangan kelas dan membandingkan metode balancing seperti SMOTE, SMOTEENN, atau SMOTETomek. Penelitian selanjutnya juga dapat membandingkan metode penelitian dengan Random Forest, Support Vector Machine, XGBoost, dan LightGBM.

4. KESIMPULAN

Penelitian ini mengoptimasi algoritma K-Nearest Neighbor untuk klasifikasi Penyakit Ginjal Kronis pada data tidak seimbang menggunakan seleksi fitur Mutual Information dan GridSearchCV. Dataset terdiri atas 400 sampel dengan 250 sampel CKD dan 150 sampel not CKD. Hasil pengujian hold-out menunjukkan bahwa KNN + GridSearchCV memperoleh performa terbaik dengan accuracy, precision, recall, F1-score, dan AUC sebesar 100,00%. Selain itu, evaluasi menggunakan 10-Fold Cross Validation menunjukkan bahwa seluruh 400 data telah digunakan sebagai data uji secara bergantian, dengan rata-rata akurasi tertinggi sebesar 99,25% pada skenario KNN + GridSearchCV. Skenario KNN + Feature Selection + GridSearchCV memperoleh accuracy 98,75%, recall 100,00%, dan F1-score 99,01%, sehingga tetap relevan karena menggunakan subset fitur yang lebih ringkas dan mendukung interpretabilitas model. Berdasarkan hasil pengujian pada dataset yang digunakan, penerapan GridSearchCV memberikan peningkatan performa dibandingkan skenario tanpa optimasi hyperparameter. Namun demikian, hasil tersebut masih terbatas pada dataset penelitian ini sehingga diperlukan validasi lebih lanjut menggunakan dataset yang lebih besar dan beragam untuk menilai kemampuan generalisasi model. Penelitian selanjutnya disarankan menggunakan dataset klinis yang lebih besar dan membandingkan metode balancing untuk meningkatkan generalisasi model pada data medis tidak seimbang. Oleh karena itu, penelitian selanjutnya disarankan menggunakan dataset yang lebih besar dan lebih beragam serta melakukan validasi eksternal agar kemampuan generalisasi model dapat dievaluasi secara lebih komprehensif.

REFERENCES

- [1] I. Wisnuadji Gamadarenda and I. Waspada, "Implementasi Data Mining Untuk Deteksi Penyakit Ginjal Kronis (PGK) Menggunakan K-Nearest Neighbor (KNN) Dengan Backward Elimination," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 7, no. 2, pp. 417–426, 2020, doi: 10.25126/jtiik.202071896.
- [2] D. Anggraini, "Aspek Klinis Dan Pemeriksaan Laboratorium Penyakit Ginjal Kronik," *An-Nadaa Jurnal Kesehatan Masyarakat*, vol. 9, no. 2, p. 236, 2022, doi: 10.31602/ann.v9i2.9229.
- [3] R. Dewi and A. Mustofa, "Penurunan Intensitas Rasa Haus Pasien Penyakit Ginjal Kronik Yang Menjalani Hemodialisa Dengan Menghisap Es Batu," *Ners Muda*, vol. 2, no. 2, p. 17, 2021, doi: 10.26714/nm.v2i2.7154.
- [4] R. Simorangkir, T. M. Andayani, and C. Wiedyaningsih, "Faktor-Faktor yang Berhubungan dengan Kualitas Hidup Pasien Penyakit Ginjal Kronis yang Menjalani Hemodialisis," *Jurnal Farmasi Dan Ilmu Kefarmasian Indonesia*, vol. 8, no. 1, p. 83, 2021, doi: 10.20473/jfiki.v8i12021.83-90.
- [5] N. P. Nugraha, R. Azim, S. Z. Daffa, and P. S. Ningayu, "Perbandingan Akurasi Metode Naive Bayes dan Metode KNN untuk Memprediksi Gagal Ginjal Kronis," *Jurnal Rekayasa Elektro Sriwijaya*, vol. 5, no. 1, pp. 1–10, 2023, doi: 10.36706/jres.v5i1.63.
- [6] H. D. Siswaja and Y. Ramdhani, "Pendekatan Algoritma Neural Network dan Genetic Algorithm Untuk Prediksi Penyakit Ginjal Kronis," *Jurnal Responsif*, vol. 6, no. 2, pp. 232–239, 2024.

- [7] U. Hasanah, N. R. Dewi, L. Ludiana, A. T. Pakarti, and A. Inayati, “Analisis Faktor-Faktor Risiko Terjadinya Penyakit Ginjal Kronik Pada Pasien Hemodialisis,” *Jurnal Wacana Kesehatan*, vol. 8, no. 2, p. 96, 2023, doi: 10.52822/jwk.v8i2.531.
- [8] M. Rizal, M. Z. Syahaf, S. R. Priyambodo, and Y. Rhamdani, “Optimasi Algoritma Naive Bayes Menggunakan Forward Selection Untuk Klasifikasi Penyakit Ginjal Kronis,” *Naratif*, vol. 5, no. 1, pp. 71–80, 2023, doi: 10.53580/naratif.v5i1.200.
- [9] Q. A’yuniyah et al., “Implementasi Algoritma Naive Bayes Classifier untuk Klasifikasi Penyakit Ginjal Kronik,” *Jurnal Sistem Komputer dan Informatika*, vol. 4, no. 1, p. 72, 2022, doi: 10.30865/json.v4i1.4781.
- [10] G. N. Ahmad, H. Fatima, S. Ullah, A. Salah Saidi, and Imdadullah, “Efficient Medical Diagnosis of Human Heart Diseases Using Machine Learning Techniques With and Without GridSearchCV,” *IEEE Access*, vol. 10, pp. 80151–80173, 2022, doi: 10.1109/ACCESS.2022.3165792.
- [11] V. Wulandari, W. J. Sari, Z. Alfian, L. Legito, and T. Arifianto, “Implementasi Algoritma Naive Bayes Classifier dan K-Nearest Neighbor untuk Klasifikasi Penyakit Ginjal Kronik,” *MALCOM*, vol. 4, no. 2, pp. 710–718, 2024, doi: 10.57152/malcom.v4i2.1229.
- [12] N. Fatimah Indrianti, A. Kania Ningsih, and R. Ilyas, “Implementasi Data Mining Untuk Klasifikasi Penyakit Gagal Ginjal Kronis Menggunakan Metode K-Nearest Neighbor,” *JATI*, vol. 8, no. 2, pp. 2255–2260, 2024, doi: 10.36040/jati.v8i2.9464.
- [13] E. Laksono, A. Basuki, and F. Bachtiar, “Optimization of K Value in KNN Algorithm for Spam and Ham Email Classification,” *Jurnal RESTI*, vol. 4, no. 2, pp. 377–383, 2020, doi: 10.29207/resti.v4i2.1845.
- [14] S. T. Kusuma and T. B. Sasongko, “Optimasi K-Nearest Neighbor dengan Grid Search CV pada Prediksi Kanker Paru-Paru,” *The Indonesian Journal of Computer Science*, vol. 12, no. 4, pp. 2162–2171, 2023, doi: 10.33022/ijcs.v12i4.3267.
- [15] R. Fadilah, Y. H. Chrisnanto, and G. Abdillah, “Pengaruh metode pengukuran jarak dan SMOTE pada klasifikasi penilaian kredit,” *JIRE*, vol. 7, no. 2, pp. 193–202, 2024.
- [16] N. Rikatsih, M. Anshori, R. Siwi Pradini, and F. Faurika, “K-Nearest Neighbor Method for Early Detection of Diabetes Patients Based on Symptoms and Clinical Data,” *Inform*, vol. 9, no. 2, pp. 187–193, 2024, doi: 10.25139/inform.v9i2.8582.
- [17] Y. Putra Dinata, M. Fikry, F. Yanto, and E. Pandu Cynthia, “Analisis Sentimen Terhadap Sebuah Figur Publik di Twitter Menggunakan Metode K-Nearest Neighbor,” *KLIK*, vol. 4, no. 6, pp. 2822–2829, 2024, doi: 10.30865/klik.v4i6.1904.
- [18] Q. A’yuniyah et al., “Implementasi Algoritma Naive Bayes Classifier (NBC) untuk Klasifikasi Penyakit Ginjal Kronik,” *Jurnal Sistem Komputer dan Informatika*, vol. 4, no. 1, p. 72, 2022, doi: 10.30865/json.v4i1.4781.
- [19] E. Saputro and D. Rosiyadi, “Penerapan Metode Random Over-Under Sampling Pada Algoritma Klasifikasi Penentuan Penyakit Diabetes,” *Bianglala Informatika*, vol. 10, no. 1, pp. 42–47, 2022, doi: 10.31294/bi.v10i1.11739.
- [20] K. Widya Kayohana, “Klasifikasi Penyakit Hati Menggunakan Random Forest Dan KNN,” *JATI*, vol. 8, no. 4, pp. 7924–7929, 2024, doi: 10.36040/jati.v8i4.10457.
- [21] R. Rozaq, “Klasifikasi Penyakit Dengue Menggunakan Algoritma K-Nearest Neighbors Berbasis Flask,” *Remik*, vol. 6, no. 3, pp. 359–369, 2022, doi: 10.33395/remik.v6i3.11501.
- [22] N. M. Putry, “Komparasi Algoritma KNN Dan Naive Bayes Untuk Klasifikasi Diagnosis Penyakit Diabetes Mellitus,” *EVOLUSI*, vol. 10, no. 1, pp. 45–57, 2022, doi: 10.31294/evolusi.v10i1.12514.
- [23] I. P. Putri, “Analisis Performa Metode K-Nearest Neighbor (KNN) dan Crossvalidation pada Data Penyakit Cardiovascular,” *Indonesian Journal of Data and Science*, vol. 2, no. 1, pp. 21–28, 2021, doi: 10.33096/ijodas.v2i1.25.