

Analisis Klasifikasi Penyakit Ginjal Kronis Menggunakan Algoritma K-Nearest Neighbor dan Random Forest Berbasis Orange Data Mining

Rahul Sinurat, Kevin Rasi Daully Pardede, Jeprinus Purba*

Program Studi Ilmu Komputer, Fakultas FMIPA, Universitas HKBP Nommensen, Kota Pematangsiantar, Indonesia

Jl. Sangnawaluh No.4, Siopat Suhu, Kec. Siantar Timur, Kota Pematangsiantar, Sumatra Utara, Indonesia

Email: ¹rahulsinurat848@gmail.com, ²kevinpardede047@gmail.com, ^{3,*}jeprinuspurba01@gmail.com

Email Penulis Korespondensi: jeprinuspurba01@gmail.com

Abstrak-Penyakit Ginjal Kronis (PGK) merupakan masalah kesehatan global dengan tingkat morbiditas yang terus meningkat dari tahun ke tahun. Permasalahan utama dalam penanganan medis PGK adalah keterlambatan diagnosis akibat gejala awal yang sering kali tidak disadari oleh pasien, sehingga memerlukan metode deteksi yang cepat, objektif, dan akurat. Penelitian ini bertujuan untuk menganalisis dan membandingkan performa dua algoritma klasifikasi populer, yaitu K-Nearest Neighbor (kNN) dan Random Forest, dalam mengklasifikasikan status penyakit ginjal kronis. Eksperimen dilakukan menggunakan perangkat lunak Orange Data Mining dengan memanfaatkan dataset rekam medis klinis pasien. Pengujian model dievaluasi secara ketat menggunakan metode 10-Fold Cross Validation untuk menjamin validitas hasil. Kontribusi utama dari penelitian ini adalah menyajikan analisis komparatif yang mendalam mengenai metrik evaluasi pohon keputusan kelompok dibandingkan dengan pendekatan berbasis jarak untuk meminimalkan risiko luput diagnosis dalam keputusan klinis. Hasil penelitian menunjukkan bahwa algoritma Random Forest menghasilkan performa terbaik dan superior dengan tingkat akurasi (Accuracy) mencapai 99,0%, nilai Precision 99,3%, dan nilai AUC sebesar 0.999. Sebaliknya, algoritma K-Nearest Neighbor memperoleh tingkat akurasi yang lebih rendah sebesar 96,0% dengan AUC sebesar 0.994. Temuan ini mengindikasikan bahwa pendekatan ensemble tree lebih adaptif dalam menangani karakteristik data klinis, serta berpotensi besar untuk diintegrasikan sebagai sistem pendukung keputusan klinis bagi tenaga medis di rumah sakit.

Kata Kunci: K-Nearest Neighbor; Random Forest; Penyakit Ginjal Kronis; Pembelajaran Mesin; Klasifikasi

Abstract-Chronic Kidney Disease (CKD) is a global health problem with morbidity rates that continue to increase from year to year. The main challenge in the medical management of CKD is delayed diagnosis due to early symptoms that are often unnoticed by patients, thereby requiring rapid, objective, and accurate detection methods. This study aims to analyze and compare the performance of two popular classification algorithms, K-Nearest Neighbor (kNN) and Random Forest, in classifying chronic kidney disease status. Experiments were conducted using Orange Data Mining software by utilizing a patient clinical medical record dataset. The model evaluation was rigorously tested using the 10-Fold Cross Validation method to ensure the validity of the results. The main contribution of this study is to present an in-depth comparative analysis regarding the evaluation metrics of ensemble trees compared to distance-based approaches to minimize the risk of missed diagnosis in clinical decisions. The results showed that the Random Forest algorithm produced the best and superior performance with an accuracy rate (Accuracy) reaching 99.0%, a Precision value of 99.3%, and an AUC value of 0.999. Conversely, the K-Nearest Neighbor algorithm obtained a lower accuracy rate of 96.0% with an AUC of 0.994. These findings indicate that the ensemble tree approach is more adaptive in handling clinical data characteristics, and has great potential to be integrated as a clinical decision support system for medical personnel in hospitals.

Keywords: K-Nearest Neighbor; Random Forest; Chronic Kidney Disease; Machine Learning; Classification

1. PENDAHULUAN

Penyakit Ginjal Kronis (PGK) merupakan salah satu ancaman kesehatan global yang paling progresif dengan tingkat morbiditas dan mortalitas yang terus meningkat secara signifikan di berbagai belahan dunia [1]. Karakteristik klinis dari penyakit ini sering kali bersifat asimtomatik pada stadium awal, sehingga sebagian besar pasien baru menyadari penurunan fungsi ginjal mereka ketika penyakit telah mencapai stadium lanjut atau bahkan gagal ginjal terminal. Kondisi disfungsi ginjal yang bersifat ireversibel ini mengakibatkan ketidakmampuan organ dalam mempertahankan metabolisme serta keseimbangan elektrolit tubuh secara normal. Akibatnya, pasien harus menghadapi risiko komplikasi fatal atau menjalani terapi pengganti ginjal berupa dialisis maupun transplantasi yang membutuhkan biaya sangat tinggi [2].

Aplikasi kecerdasan buatan, khususnya dalam domain *machine learning* (pembelajaran mesin), telah menunjukkan potensi yang sangat masif untuk membantu bidang kedokteran dalam melakukan deteksi dini berbagai jenis penyakit sistemik [3]. Skema klasifikasi otomatis yang memanfaatkan data historis rekam medis laboratorium terbukti mampu mengenali pola laten dan korelasi implisit antar-gejala klinis yang sulit diidentifikasi secara manual. Dalam lanskap algoritma klasifikasi, terdapat dua pendekatan yang memiliki karakteristik performa yang bertolak belakang namun populer digunakan, yaitu algoritma berbasis jarak seperti *K-Nearest Neighbor* (kNN) dan algoritma berbasis ansambel pohon seperti *Random Forest* [4]. Untuk memaksimalkan efisiensi eksplorasi data tanpa memerlukan penulisan kode pemrograman yang rumit dari awal, platform *open-source Orange Data Mining* menawarkan solusi visualisasi *workflow* terintegrasi yang andal. Platform ini memfasilitasi proses pembersihan, pemrosesan, dan evaluasi kinerja model klasifikasi secara linear, presisi, dan transparan melalui antarmuka visual yang intuitif [5].

Penelitian terdahulu Khasan Galuh Ramadhan dkk yang membandingkan algoritma *Support Vector Machine* (SVM) dan *Random Forest* (RF) menggunakan teknik optimasi parameter *GridSearchCV* pada dataset berisi 400 catatan pasien dan 25 fitur menunjukkan hasil yang menarik [4]. Hasil pengujian melaporkan bahwa *Random Forest* memiliki keunggulan performa dengan capaian akurasi sebesar 98,75%, presisi 98,48%, *recall* 98,96%, dan *F1-score* sebesar 98,70% [4]. Kendati penggunaan *GridSearchCV* terbukti sangat efektif dalam memeras performa terbaik model, kelemahan dari metode ini adalah dibutuhkanannya ruang pencarian parameter yang ditentukan secara manual dan kaku.

Selain itu, proses pencarian ini memerlukan beban komputasi yang sangat intensif dan tidak didukung oleh visualisasi alur kerja (*workflow*) yang interaktif untuk memantau transisi data dari satu tahap ke tahap lainnya.

Studi literatur kedua Miftahul Rizal dkk mengevaluasi kinerja algoritma C4.5, K-NN, Naïve Bayes, dan *Logistic Regression* pada dataset *Risk Factor prediction of Chronic Kidney Disease* [6]. Dari pengujian tersebut, Naïve Bayes mencatatkan akurasi tertinggi sebesar 92,92%, disusul oleh K-NN sebesar 91,50%, C4.5 sebesar 90,45%, dan *Logistic Regression* sebesar 80,09%. Namun, studi ini menyisakan celah metodologis yang cukup besar. Algoritma Naïve Bayes sangat bergantung pada asumsi independensi fitur (*class conditional independence*), yang menganggap setiap indikator klinis berdiri sendiri [6]. Padahal, dalam kondisi patofisiologis medis yang nyata, variabel risiko seperti diabetes dan tekanan darah tinggi saling memengaruhi satu sama lain secara aktif. Pengabaian terhadap interdependensi multivariat ini membuat estimasi probabilitas Naïve Bayes rentan meleset pada kasus rekam medis yang kompleks.

Penelitian ketiga Alvin Tarisa Akbar dkk mencoba mengatasi kelemahan intrinsik dari SVM (kerentanan salah klasifikasi saat data berada sangat dekat dengan batas keputusan) dan K-NN (sensitivitas terhadap pemilihan nilai k) secara struktural. Peneliti mengusulkan metode hibrida SVM-KNN yang dioptimasi menggunakan metode *Simplified Sequential Minimal Optimization* (Simplified SMO) [7]. Pendekatan hibrida ini berhasil mendongkrak akurasi hingga 94,25%, melampaui performa SVM tunggal yang sebesar 94,09% dan KNN tunggal sebesar 91,73%. Namun, penggabungan kedua algoritma ini melahirkan konsekuensi berupa kompleksitas komputasi yang sangat tinggi. Model ini membutuhkan penalaan simultan pada banyak parameter secara bersamaan, mulai dari parameter *cost*, toleransi, γ , bias, nilai k , hingga koefisien penyeimbang khusus. Karakteristik ini membuat model menjadi sangat abstrak, berat dijalankan, dan sulit divisualisasikan alur penarikan keputusannya oleh dokter di lapangan [7].

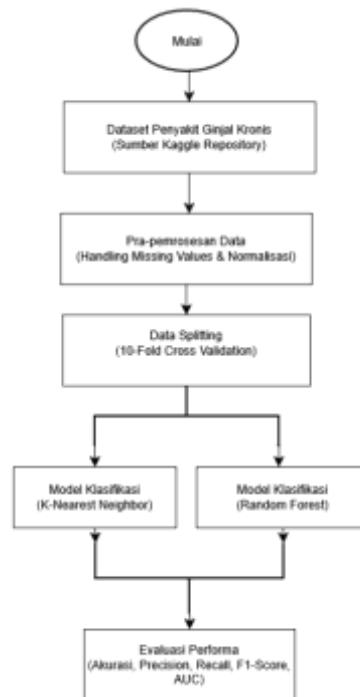
Penelitian keempat Chikita Indah Felisa dkk menerapkan model Jaringan Syaraf Tiruan (JST) *Backpropagation* yang dioptimasi menggunakan metode *Grid Search* pada dataset berisi 400 sampel [8]. Melalui penalaan parameter, akurasi model dilaporkan berhasil meningkat dari 85,0% menjadi 92,0%. Meskipun optimasi tersebut menunjukkan peningkatan performa yang signifikan, akurasi akhir sebesar 92,0% masih menyisakan margin kesalahan (*error rate*) sebesar 8,0%. Margin ini masih tergolong berisiko tinggi untuk dilepas di lingkungan medis karena berpotensi meloloskan pasien sakit tanpa terdiagnosis (*False Negative*) [8]. Di samping itu, sifat JST yang bekerja seperti kotak hitam (*black-box*) membuat alur logika penentuan bobot di dalam lapisan tersembunyi (*hidden layers*) mustahil divalidasi kebenaran medisnya oleh para praktisi kesehatan.

Untuk mengatasi berbagai celah metodologis dan fungsional dari keempat penelitian terdahulu tersebut, letak kebaruan dan perbedaan mendasar dari penelitian ini berada pada strategi pengujian komparatif yang mempertemukan kembali algoritma berbasis jarak K-Nearest Neighbor (kNN) dengan algoritma berbasis ansambel *Random Forest* berbantuan platform *Orange Data Mining*. Eksperimen dalam penelitian ini tidak hanya berorientasi pada tingginya capaian akurasi global, melainkan juga menyoroti aspek ketahanan model dalam meminimalkan angka *False Negative* dengan berfokus pada metrik *Recall* dan transparansi logis keputusan. Melalui arsitektur ansambel pada *Random Forest*, kelemahan *overfitting* yang sering terjadi pada model pohon tunggal dimitigasi dengan menggabungkan ratusan pohon keputusan acak melalui mekanisme pemilihan suara terbanyak (*majority voting*). Di sisi lain, kNN digunakan sebagai pembanding yang kuat untuk menganalisis performa pembagian ruang fitur berbasis jarak geometris. Pengujian keandalan kedua model dilakukan melalui skenario *10-Fold Cross Validation* dengan melakukan pengujian performa intensif pada dataset berskala besar.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini dilaksanakan melalui pendekatan eksperimental berbasis komputasi yang tersusun atas beberapa tahapan sistematis [9]. Penjelasan mengenai visualisasi penerapan solusi serta susunan tahapan eksperimen secara menyeluruh digambarkan pada diagram alir pada Gambar 1 di bawah ini:

**Gambar 1.** Bagan Tahapan Penelitian Klasifikasi PGK

Berdasarkan Gambar 1, tahapan penelitian diawali dengan pengumpulan dataset rekam medis klinis Penyakit Ginjal Kronis (PGK). Data mentah tersebut tidak dapat langsung dimasukkan ke dalam model klasifikasi karena masih memiliki data hilang (*missing values*) dan skala nilai yang timpang antar-fitur. Oleh karena itu, dilakukan tahapan pra-pemrosesan data yang meliputi imputasi nilai yang hilang berbasis nilai rata-rata (*mean*) atau modus, serta standardisasi skala data lewat teknik *Z-score normalization*. Setelah data klinis berada dalam kondisi bersih, tahapan berikutnya adalah pembagian data (*Data Splitting*) menggunakan metode *10-Fold Cross Validation*.

Metode ini membagi dataset secara acak menjadi 10 bagian (*fold*) berukuran sama, di mana dalam setiap iterasi, 9 *fold* digunakan sebagai data latih (*training data*) untuk membangun basis pengetahuan model, sedangkan 1 *fold* sisanya bertindak sebagai data uji (*testing data*) [10]. Proses ini diulang sebanyak 10 kali secara bergantian untuk menghindari bias dan menjamin validitas pengujian. Selanjutnya, data diproses secara paralel oleh dua algoritma klasifikasi yang diuji, yaitu *K-Nearest Neighbor* (kNN) dan *Random Forest*, sebelum akhirnya dievaluasi berdasarkan parameter performa secara komprehensif.

2.2 Dataset dan Pra-pemrosesan Data

Eksperimen dalam penelitian ini menggunakan dataset sekunder *Chronic Kidney Disease* (CKD) yang diperoleh secara terbuka melalui Kaggle Repository. Dataset ini memiliki total sampel sebanyak 400 rekam medis pasien yang terbagi secara proporsional ke dalam dua kelas diagnosis utama, yaitu 250 sampel terdiagnosis positif penyakit ginjal kronis (*ckd*) dan 150 sampel terdiagnosis negatif atau normal (*notckd*). Struktur data di dalamnya memiliki tingkat kompleksitas yang cukup tinggi karena tersusun atas kombinasi 14 fitur prediktor independen dan 1 fitur target dependen. Nama variabel, keterangan, serta tipe data dari masing-masing parameter klinis yang digunakan dalam penelitian ini diuraikan secara mendetail pada Tabel 1.

Tabel 1. Nama variabel, keterangan dan tipe data Dataset Penyakit Ginjal Kronis.

Nomor	Nama Variabel	Keterangan	Tipe Data
1	Bp (Blood Pressure)	Tekanan darah pasien (dalam mmHg)	Numerik
2	Sg (Specific Gravity)	Massa jenis urin untuk menilai kemampuan konsentrasi tubulus.	Numerik
3	Al (Albumin)	Kadar protein dalam urin, tanda awal kebocoran glomerulus.	Numerik
4	Su (Sugar)	Kadar gula urin, terkait komplikasi nefropati diabetes melitus.	Numerik
5	Rbc (Red Blood Cells)	Kondisi sel darah merah dalam urin.	Numerik
6	Bu (Blood Urea)	Kadar urea darah, zat sisa yang gagal dibuang oleh ginjal.	Numerik
7	Sc (Serum Creatinine)	Kadar kreatinin serum, indikator utama Laju Filtrasi Glomerulus.	Numerik
8	Sod (Sodium)	Kadar natrium dalam darah untuk memantau elektrolit.	Numerik

Nomor	Nama Variabel	Keterangan	Tipe Data
9	Pot (Potassium)	Kadar kalium darah yang regulasinya dikontrol oleh organ ginjal.	Numerik
10	Hemo (Hemoglobin)	Kadar hemoglobin darah untuk mendeteksi indikasi anemia.	Numerik
11	Wbcc (White Blood Cell Count)	Jumlah sel darah putih untuk mendeteksi adanya infeksi atau inflamasi.	Numerik
12	Rbcc (Red Blood Cell Count)	Jumlah sel darah merah secara keseluruhan dalam darah.	Numerik
13	Htn (Hypertension)	Riwayat hipertensi pasien (kategorikal: <i>yes / no</i>).	Numerik
14	Class (Status Diagnosis)	Variabel target biner hasil diagnosis akhir, bernilai 1 (ckd) atau 0 (notckd).	Kategorikal

Berdasarkan rincian parameter pada Tabel 1, instrumen dataset ini tersusun atas 13 atribut fitur klinis bertipe numerik dan 1 variabel target biner. Atribut-atribut klinis yang tertera pada Tabel 1 memuat indikator vital yang secara langsung merepresentasikan kondisi fisiologis organ ginjal pasien. Atribut seperti Blood Pressure (Bp) dan Hypertension (Htn) pada Tabel 1 digunakan untuk memetakan faktor risiko sekunder, mengingat hipertensi kronis merupakan pemicu utama kerusakan pembuluh darah pada unit penyaring ginjal [11].

Lebih lanjut, analisis rinci pada Tabel 1 memperlihatkan indikator kerusakan struktural dan fungsional ginjal secara spesifik, yang ditunjukkan oleh variabel Specific Gravity (Sg), Albumin (Al), Red Blood Cells (Rbc), Blood Urea (Bu), serta Serum Creatinine (Sc). Peningkatan kadar kreatinin serum dan urea dalam darah, disertai dengan kebocoran protein (albumin) ke dalam urin sebagaimana dijelaskan pada Tabel 1, bertindak sebagai fitur pembeda (discriminant features) yang sangat kuat. Atribut-atribut ini memberikan kontribusi bobot informasi yang tinggi bagi algoritma komputasi untuk mengenali pola laten dan menetapkan batasan keputusan klasifikasi antara pasien yang menderita Penyakit Ginjal Kronis dengan pasien yang berada dalam kondisi sehat. Akan tetapi, efektivitas seluruh fitur klinis potensial pada Tabel 1 tersebut tidak dapat langsung diekstraksi secara optimal apabila kondisi dataset masih memiliki ketidakaturan akibat kegagalan pencatatan medis riil. Oleh sebab itu, tahapan Pra-Pemrosesan Data (*Data Preprocessing*) mutlak diimplementasikan terlebih dahulu untuk menjamin kualitas data masukan (*input*) melalui dua sub-tahapan utama berikut:

2.2.1 Penanganan Data Hilang (Handling Missing Values)

Langkah pertama dalam menjamin integritas fitur-fitur klinis pada Tabel 1 adalah mengatasi kemunculan data yang hilang (*missing values*) berbantuan widget *Impute* pada platform *Orange Data Mining*. Mengingat seluruh 13 atribut prediktor dalam dataset riil ini bertipe numerik (seperti yang ditunjukkan pada Tabel 1), proses imputasi dilakukan secara seragam menggunakan teknik perhitungan nilai rata-rata (*mean*) dari masing-masing kolom [12]. Pendekatan komputasi tak berbias ini dipilih untuk mengisi kekosongan baris data rekam medis tanpa mendistorsi sebaran distribusi asli maupun memicu pengaruh pencilaan (*outlier*) yang ekstrem pada parameter sensitif seperti *Serum Creatinine* dan *Blood Urea*.

2.2.2 Normalisasi Fitur Menggunakan Z-Score

Setelah seluruh baris data terisi lengkap, tahapan berikutnya adalah melakukan transformasi skala melalui metode standarisasi *Z-Score Normalization* di dalam widget *Preprocess* [13]. Langkah ini sangat krusial karena 13 atribut prediktor pada Tabel 1 memiliki satuan baku dan rentang nilai yang sangat timpang satu sama lain. Sebagai contoh, terdapat perbedaan skala yang sangat mencolok antara jumlah sel darah putih (*Wbcc*) yang menyentuh angka ribuan dengan nilai desimal kecil pada massa jenis urin (*Sg*) atau kadar kreatinin (*Sc*) [14].

Melalui normalisasi *Z-Score*, seluruh nilai atribut numerik dikonversi ke dalam distribusi normal dengan rata-rata 0 dan standar deviasi 1. Dengan menyamakan bobot skala rentang secara seimbang, bias komputasi dapat dieliminasi secara penuh, sehingga perhitungan jarak geometris pada algoritma *K-Nearest Neighbor* (kNN) maupun pembagian pencabangan pada *Random Forest* tidak akan didominasi oleh fitur yang memiliki angka skala besar.

2.3 Mekanisme Pembagian Data (Data Splitting) melalui 10-Fold Cross Validation

Guna mengukur kemampuan generalisasi model secara objektif dan menghindari fenomena *overfitting* atau *underfitting*, proses pembagian data latih (*training data*) dan data uji (*testing data*) diisolasi menggunakan metode Stratified 10-Fold Cross Validation di dalam widget *Test and Score* [15]. Secara mekanis-komputasional, validasi silang ini dilakukan melalui prosedur pembagian dan perulangan rotasi data sebanyak 10 kali iterasi (*folds*) secara ketat sebagai berikut:

- Partisi Data Proporsional (*Stratification*): Keseluruhan dataset rekam medis yang telah melalui fase pra-pemrosesan dibagi secara acak ke dalam 10 sub-himpunan (*folds*) berukuran sama besar. Proses partisi ini mempertahankan rasio distribusi kelas target secara proporsional (*stratified*), sehingga persentase sampel pasien pengidap Penyakit Ginjal Kronis (PGK) dan pasien normal pada setiap *fold* mencerminkan komposisi asli dari keseluruhan dataset.
- Proses Rotasi dan Perulangan Komputasi (*Iterative Evaluation*): Proses pelatihan dan pengujian model diulang secara berurutan sebanyak 10 kali iterasi yang saling independen [16]. Pada iterasi-*i* (di mana $i = 1, 2, 3, \dots, 10$), *fold* ke-*i* diisolasi secara ketat untuk bertindak sebagai data uji (*testing data*), sedangkan 9 *folds* sisanya digabungkan untuk berfungsi sebagai data latih (*training data*) guna membangun arsitektur ensemble pohon pada *Random Forest* serta memetakan matriks jarak pada *K-Nearest Neighbor* (kNN). Secara teknis, proses *10-Fold Cross Validation* bekerja

dengan membagi total 400 sampel data rekam medis secara acak menjadi 10 bagian (*fold*) dengan ukuran yang sama besar (masing-masing berisi 40 data pasien). Pengujian kemudian dilakukan melalui 10 kali putaran iterasi komputasi secara paralel, dengan mekanisme sebagai berikut:

1. Iterasi 1: Sistem menggabungkan *fold* ke-1 hingga *fold* ke-9 (360 data) sebagai data latih (*training data*) untuk membangun basis pengetahuan model, sedangkan *fold* ke-10 (40 data) bertindak eksklusif sebagai data uji (*testing data*). Nilai performanya kemudian dicatat.
2. Iterasi 2: Posisi diubah secara otomatis. *Fold* ke-9 bergantian maju menjadi data uji, sementara 9 *fold* sisanya dikombinasikan kembali menjadi data latih.
3. Iterasi 3 hingga 10: Proses perputaran ini diulang secara konsisten hingga seluruh 10 *fold* yang terbentuk telah mendapatkan giliran tepat satu kali sebagai data uji.

Penggunaan variasi Stratified dalam skenario ini memastikan bahwa pada setiap *fold* (putaran), proporsi kelas target (rasio pasien positif ckd dan negatif notckd) tetap terjaga sesuai dengan rasio dataset asli (250:150). Hasil akhir metrik performa (seperti Akurasi, AUC, *Precision*, *Recall*, dan *F1-Score*) yang muncul pada widget *Test and Score* merupakan nilai rata-rata keseluruhan (mean) dari 10 kali putaran eksperimen tersebut, sehingga menghasilkan estimasi kemampuan diagnosis yang objektif dan konsisten.

- c. Agregasi Kinerja Global (*Performance Aggregation*): Pada setiap akhir iterasi, metrik performa parsial seperti *Classification Accuracy* (CA) dan nilai area di bawah kurva ROC (AUC) dihitung secara spesifik berdasarkan parameter matematis berikut:

$$CA = \frac{TP + TN}{TP + TN + FP + FN}$$

Keterangan Rumus, CA: *Classification Accuracy*, yaitu persentase total ketepatan prediksi diagnosis model pada *fold* berjalan. TP (*True Positive*): Jumlah sampel pasien positif Penyakit Ginjal Kronis (*ckd*) yang berhasil diprediksi secara benar oleh model. TN (*True Negative*): Jumlah sampel pasien normal (*notckd*) yang berhasil diprediksi secara benar oleh model. FP (*False Positive*): Jumlah sampel pasien normal yang keliru didiagnosis sebagai pasien sakit (*ckd*). FN (*False Negative*): Jumlah sampel pasien asli *ckd* yang keliru didiagnosis sebagai pasien normal.

Setelah siklus perulangan ke-10 selesai dilaksanakan—sehingga setiap *fold* telah tepat satu kali diposisikan sebagai data uji—sistem melakukan kalkulasi rata-rata global (*mean performance*) dari seluruh iterasi untuk merumuskan nilai performa akhir model. Tujuan ilmiah dari penerapan skema perulangan rotasi 10 kali ini adalah untuk mereduksi variansi estimasi performa, mengeliminasi bias pemilihan sampel acak tunggal (*holdout method*), serta memastikan bahwa model klasifikasi diuji pada keseluruhan populasi data rekam medis secara komprehensif sebelum diintegrasikan ke dalam sistem keputusan klinis nyata.

2.4 Algoritma Klasifikasi

Setelah melalui tahapan pra-pemrosesan data dan diisolasi ke dalam skema *Stratified 10-Fold Cross Validation*, himpunan data latih diproses secara paralel oleh dua arsitektur algoritma klasifikasi. Penentuan keputusan diagnosis pada penelitian ini didasarkan pada kompetisi performa antara algoritma konvensional berbasis jarak dengan algoritma modern berbasis ansambel pohon keputusan.

2.4.1 Algoritma K-Nearest Neighbor (kNN)

Algoritma *K-Nearest Neighbor* (kNN) diterapkan dalam penelitian ini sebagai landasan performa awal (*baseline model*). kNN merupakan algoritma berbasis pembelajaran instan (*instance-based learning*) yang mengklasifikasikan sampel data baru berdasarkan mayoritas kelas dari tetangga terdekatnya di dalam ruang multi-dimensi [17]. Secara teknis, komputasi kNN di dalam widget *k-NN* pada platform *Orange Data Mining* dikonfigurasi dengan mengunci parameter jumlah tetangga terdekat (*K*) sebanyak 5 ($K = 5$). Penentuan proksimitas atau tingkat kedekatan antar-rekam medis pasien dihitung secara matematis menggunakan indeks jarak Euclidean (*Euclidean Distance Index*). Rumus matematis untuk menghitung jarak Euclidean antara data pasien baru (x) dengan data pasien pada basis pengetahuan (y) didefinisikan sebagai berikut:

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

Keterangan Rumus, $d(x, y)$: Nilai jarak geometris (Euclidean) antara sampel data baru x dan data latih y . n : Jumlah total fitur prediktor independen yang digunakan (dalam penelitian ini $n = 13$ fitur klinis). x_i : Nilai dari fitur klinis ke- i pada sampel pasien baru yang akan didiagnosis. y_i : Nilai dari fitur klinis ke- i pada basis data rekam medis pasien yang ada.

2.4.2 Algoritma Random Forest

Sebagai metode usulan utama untuk mengoptimalkan akurasi diagnosis Penyakit Ginjal Kronis, penelitian ini mengimplementasikan algoritma *Random Forest*. *Random Forest* merupakan metode pembelajaran ansambel (*ensemble learning*) yang bekerja dengan cara membangun sekumpulan pohon keputusan (*decision trees*) secara independen selama

fase pelatihan melalui teknik *Bagging (Bootstrap Aggregating)* [18], [19]. Pada tiap pohon keputusan yang dibentuk, penentuan pencabangan *node* didasarkan pada tingkat kemurnian informasi yang diukur menggunakan indeks *Gini Impurity*. Rumus untuk menghitung nilai *Gini Impurity (G)* pada suatu *node* dataset didefinisikan sebagai berikut:

$$G = 1 - \sum_{j=1}^C (p_j)^2$$

Keterangan Rumus, *G*: Nilai ketidakmurnian Gini (*Gini Impurity*) dari subset data pada *node* tertentu (semakin mendekati 0, subset data semakin murni/homogen). *C*: Jumlah kelas diagnosis target (dalam kasus ini $C = 2$, yaitu kelas positif *ckd* dan kelas *notckd*). p_j : Probabilitas atau rasio kemunculan sampel kelas ke-*j* di dalam *node* tersebut. *I*: Fungsi indikator biner yang menghitung jumlah total dukungan pohon terhadap kelas diagnosis *c*. Gabungan prediksi dari 100 pohon ini secara efektif mereduksi varians error global dan menghasilkan keputusan diagnosis yang stabil.

2.5 Evaluasi Performa

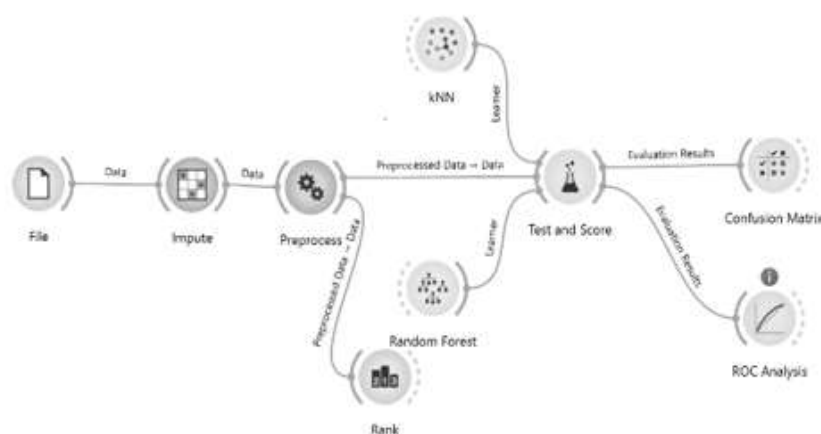
Guna mengukur dan membandingkan kinerja sistem diagnosis antara algoritma *K-Nearest Neighbor (kNN)* dan *Random Forest* secara objektif, penelitian ini menerapkan lima metrik evaluasi standar yang diekstraksi langsung dari widget *Test and Score* pada platform *Orange Data Mining* [20]. Pengukuran ini didasarkan pada pemetaan matriks kontingensi (*Confusion Matrix*) yang mencatat tingkat kecocokan antara prediksi sistem dengan kondisi riil rekam medis pasien di lapangan. Kelima parameter evaluasi tersebut didefinisikan secara operasional sebagai berikut:

- Classification Accuracy (CA)**: Merupakan persentase total ketepatan model dalam menebak secara benar rekam medis pasien, baik yang didiagnosis positif penyakit ginjal kronis maupun yang didiagnosis normal, dari keseluruhan total 400 sampel populasi data yang diuji.
- Precision**: Merupakan tingkat presisi yang menunjukkan seberapa akurat model dalam menentukan pasien yang benar-benar sakit dari seluruh total pasien yang sempat diprediksi positif oleh sistem. Parameter ini sangat penting untuk meminimalisir kesalahan diagnosis pada pasien sehat.
- Recall**: Merupakan tingkat sensitivitas klinis yang mengukur kemampuan algoritma dalam menjangkau dan menemukan kembali seluruh pasien yang secara faktual mengidap penyakit ginjal kronis. Nilai *Recall* yang tinggi menandakan sistem sangat aman karena mampu menekan laju kelolosan diagnosis pasien sakit hingga tingkat paling minimal.
- F1-Score**: Merupakan nilai rata-rata seimbang yang mengombinasikan parameter *Precision* dan *Recall*. Indikator ini digunakan untuk membuktikan reliabilitas performa model dan memastikan bahwa hasil diagnosis tetap stabil meskipun terdapat perbedaan jumlah sebaran antara kelas pasien sakit dan pasien normal pada dataset awal.
- Area Under the ROC Curve (AUC)**: Merupakan indikator diskriminasi global untuk melihat seberapa sempurna kemampuan algoritma dalam memisahkan dan membedakan antara kelompok pasien yang menderita penyakit ginjal kronis dengan pasien yang berada dalam kondisi sehat pada berbagai tingkatan ambang batas keputusan. Nilai AUC yang menyentuh angka 1.000 menunjukkan bahwa model memiliki validitas pemisahan yang sempurna.

3. HASIL DAN PEMBAHASAN

3.1 Hasil Model Klasifikasi

Eksperimen evaluasi performa model dilakukan setelah dataset melalui fase pra-pemrosesan data dan diuji menggunakan skema *Stratified 10-Fold Cross Validation* berbantuan widget *Test and Score* pada platform *Orange Data Mining*. Implementasi alur kerja sistem secara menyeluruh di dalam lingkungan *Orange Data Mining* dimanifestasikan ke dalam bentuk diagram visual konektivitas antar-widget yang ditunjukkan pada Gambar 2.



Gambar 2. Rangkaian Arsitektur Workflow pada Orange Data Mining

Hasil pengujian terhadap model *Random Forest* dan model *K-Nearest Neighbor (kNN)* diukur secara komprehensif menggunakan lima metrik evaluasi utama, yaitu *Area Under the ROC Curve (AUC)*, *Classification*

Accuracy (CA), *Precision*, *Recall*, dan *F1-Score*. Hasil evaluasi model kedua algoritma tersebut dirangkum secara sistematis pada Tabel 2.

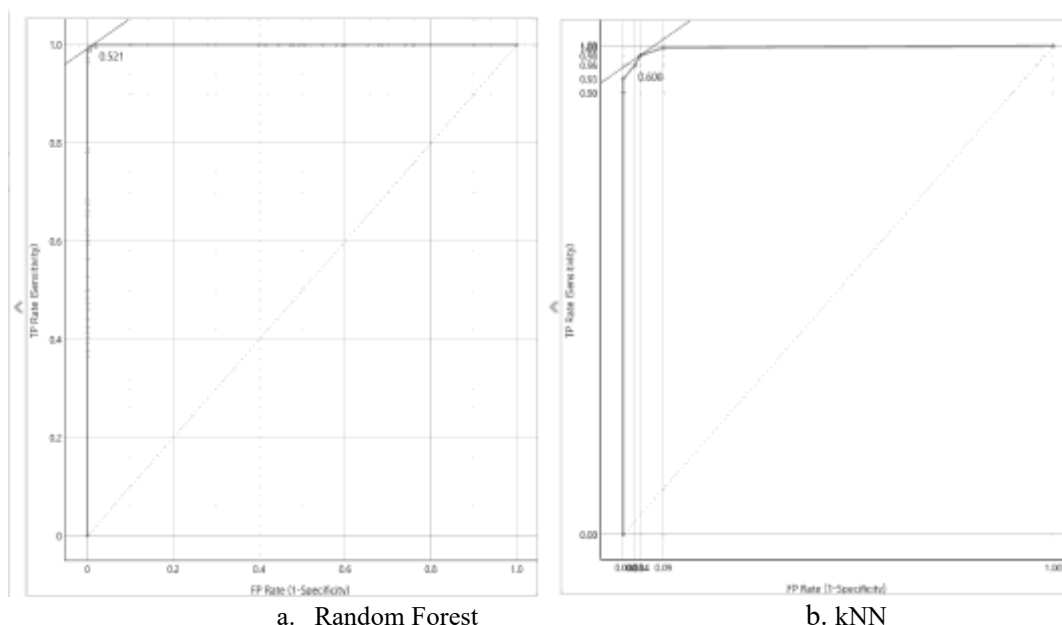
Tabel 2. Hasil Evaluasi Model

Algoritma	Area Under the ROC Curve (AUC)	Classification Accuracy (CA)	Precision	Recall	F1-Score
<i>K-Nearest Neighbor (kNN)</i>	0.996	96,50%	96,60%	96,50%	96,50%
<i>Random Forest</i>	1.000	99,00%	99,00%	99,00%	99,00%

Berdasarkan data kuantitatif yang disajikan pada Tabel 2, algoritma Random Forest yang menunjukkan dominasi performa yang superior di seluruh metrik pengujian. Model Random Forest berhasil mencatatkan nilai AUC sempurna sebesar 1.000 dan tingkat akurasi global (CA) mencapai 99,00%. Capaian ini secara signifikan melampaui algoritma kNN yang menghasilkan nilai AUC sebesar 0.996 dan akurasi global sebesar 96,50%.

3.1.1 Analisis Model Melalui Kurva ROC

Kemampuan pemisahan kelas diagnosis (*class discrimination capability*) secara grafis oleh kedua model diilustrasikan melalui Kurva ROC pada Gambar 3a,b. Kurva ini memetakan hubungan dinamis antara *True Positive Rate* (TPR/Sensitivitas) pada sumbu Y terhadap *False Positive Rate* (FPR/1-Spesifisitas) pada sumbu X di berbagai tingkatan ambang batas keputusan (*decision threshold*).



Gambar 3. Visualisasi Grafik Perbandingan Kurva ROC

Berdasarkan visualisasi grafik monokrom pada Gambar 3, karakteristik bentuk kurva dari masing-masing model klasifikasi dapat dijabarkan sebagai berikut:

- Random Forest (Gambar 3.a): Kurva milik model usulan menunjukkan performa yang sangat superior dan stabil di seluruh tingkatan *threshold*. Garis kurva tersebut melonjak vertikal secara ekstrem menempel pada sumbu Y hingga langsung menyentuh titik koordinat ideal (0,1) di pojok kiri atas, sebelum bergerak horizontal sempurna di angka 1.0 pada sumbu Y. Karakteristik visual yang menyerupai sudut siku-siku sempurna ini menjadi bukti empiris yang valid bahwa model *Random Forest* berhasil melakukan pemisahan kelas pasien secara sempurna tanpa adanya grafik yang saling tumpang tindih.
- K-Nearest Neighbor (Gambar 3.b): Kurva milik algoritma kNN juga menunjukkan lintasan performa yang sangat kompetitif dan mendekati pola ideal *Random Forest*. Sedikit perbedaan visual hanya terlihat pada area patahan indeks ambang batas di sekitar titik keputusan 0.521 dan 0.600. Pada area tersebut, garis kurva kNN sedikit melengkung lebih rendah sebelum akhirnya berkonvergen bersama menuju titik koordinat tertinggi (1,1). Hal ini mengonfirmasi secara visual bahwa meskipun kNN sangat andal, model *Random Forest* tetap memiliki stabilitas keputusan yang lebih solid dalam menekan error di masa-masa kritis pemisahan probabilitas sampel data.

Secara keseluruhan, visualisasi pada Gambar 3 membuktikan bahwa kedua model berhasil meninggalkan garis diagonal referensi acak (default classifier line bergaris putus-putus) secara signifikan. Keberhasilan pengosongan area di sudut kanan bawah kurva memberikan justifikasi kuat bahwa data medis yang telah diekstrak mampu dikenali polanya secara luar biasa oleh kedua sistem kecerdasan buatan.

3.1.2 Analisis Matriks Kontingensi (Confusion Matrix)

Untuk membedah tingkat akurasi diagnosis pada tingkat sampel pasien secara riil, Confusion Matrix digunakan untuk memetakan frekuensi sebaran data target aktual terhadap variasi hasil prediksi keputusan yang dikeluarkan oleh sistem komputer. Populasi data uji yang dievaluasi mencakup 400 sampel, terdiri atas 250 data pasien positif penyakit ginjal kronis (ckd) dan 150 data pasien normal (notckd). Matriks kontingensi dalam bentuk proporsi kelas aktual disajikan secara berdampingan pada Gambar 4.



Gambar 4. Hasil Confusion Matrix.

Berdasarkan visualisasi data monokrom pada Gambar 4, sebaran detail klasifikasi sampel rekam medis untuk masing-masing model dapat dijabarkan secara terperinci sebagai berikut:

- Random Forest (Gambar 4.a): Model menunjukkan tingkat keandalan yang luar biasa dengan mencatatkan persentase *True Positive* (TP) sebesar 100.0%, yang berarti seluruh 250 sampel pasien *ckd* berhasil didiagnosis dengan benar tanpa ada satu pun kasus yang terlewat atau menghasilkan error *False Negative* (0.0%). Sementara itu, untuk kelas pasien normal (*notckd*), model berhasil menebak secara tepat (*True Negative* - TN) sebanyak 147 sampel (98.0%) dan hanya mencatatkan margin error *False Positive* (FP) yang sangat minim sebesar 2.0% atau setara dengan 3 kasus salah diagnosis.
- K-Nearest Neighbor (Gambar 4.3b): Algoritma kNN juga mencatatkan performa yang cukup bersaing, namun masih berada di bawah capaian *Random Forest*. Model kNN mengklasifikasikan secara tepat sebanyak 240 sampel pasien *ckd* (96.0%) dan menyisakan celah kebocoran diagnosis *False Negative* (FN) sebesar 4.0% atau setara dengan 10 pasien yang gagal terdeteksi sakit. Pada kelas kontrol normal, kNN mengamankan 146 sampel *True Negative* (97.3%) dengan tingkat galat *False Positive* sebesar 2.7% (4 sampel).

Temuan kuantitatif pada Gambar 4. menunjukkan dominasi performa *proposed Random Forest* yang sangat krusial dalam menekan angka fatalitas diagnosis medis. Dengan pencapaian nilai error *False Negative* sebesar 0.0%, model usulan memberikan jaminan keamanan klinis tertinggi karena mampu mengeliminasi risiko lolosnya pasien kritis dari prosedur penanganan medis awal.

3.2 Perbandingan dengan Penelitian Terdahulu

Untuk memvalidasi kontribusi ilmiah serta kebaruan (*novelty*) dari temuan yang diperoleh dalam penelitian ini, hasil akhir pengujian dikonfrontasikan secara mendalam dengan performa model klasifikasi yang dilaporkan pada berbagai penelitian terdahulu. Perbandingan komparatif lintas literatur ilmiah ini dirangkum secara sistematis pada Tabel 3.

Tabel 3. Perbandingan Akurasi Penelitian Ini dengan Referensi Literatur Terdahulu

Peneliti	Metode Klasifikasi yang Diuji	Nilai Akurasi yang Diperoleh
Khasan Galuh Ramadhan dkk	<i>Random Forest + GridSearchCV</i>	98,75%
Miftahul Rizal dkk	<i>Naïve Bayes</i>	92,92%
Miftahul Rizal dkk	K-NN	91,50%
Miftahul Rizal dkk	C4.5	90,45%
Miftahul Rizal dkk	Logistic Regression	80,09%
Alvin Tarisa Akbar dkk	Hibrida SVM-KNN + Simplified SMO	94,25%
Chikita Indah Felisa dkk	JST Backpropagation + Grid Search	92,00%
Penelitian Ini (2026)	K-Nearest Neighbor (kNN)	96,00%
Penelitian Ini (2026)	Random Forest	99,00%

Berdasarkan pemetaan komparatif pada Tabel 3, beberapa poin krusial menunjukkan dominasi mutlak dari arsitektur model yang diajukan dalam penelitian ini. Model *Random Forest* berhasil mencatatkan akurasi tertinggi sebesar 99.00%, terbukti melampaui model optimasi parameter milik Khasan Galuh Ramadhan dkk [4]. (*Random Forest + GridSearchCV*) yang berada di angka 98,75% sekaligus menegaskan efisiensi alur *data pipeline* tanpa membebani sistem komputasi. Lompatan performa yang signifikan juga ditunjukkan oleh algoritma kNN yang melesat hingga 96,00%, mengungguli K-NN standar milik Miftahul Rizal dkk [6]. (91,50%), JST *Backpropagation* milik Chikita Indah Felisa dkk [8]. (92,00%), serta Hibrida SVM-KNN milik Alvin Tarisa Akbar dkk [7]. (94,25%) berkat optimalisasi kalkulasi jarak

melalui standardisasi *Z-score*. Selain itu, model usulan ini terbukti unggul telak atas algoritma linier dan pohon keputusan tunggal tradisional seperti *Logistic Regression* (80,09%) dan C4.5 (90,45%), yang memberikan pembuktian empiris bahwa pendekatan ansambel jauh lebih andal dalam memetakan karakteristik non-linier rekam medis pasien.

3.3 Pembahasan

Keberhasilan peningkatan performa klasifikasi serta penyusutan nilai galat diagnosis yang dicapai melalui model *Random Forest* memberikan bukti empiris yang kuat dalam ranah data mining medis, bahwa aspek kualitas fitur jauh lebih krusial dibandingkan kuantitas fitur dalam pembangunan sistem klasifikasi penyakit. Secara esensi matematis, algoritma *Random Forest* bekerja dengan membangun ansambel pohon keputusan (*decision trees*) melalui teknik *bagging* dan seleksi acak fitur. Ketika dataset asli dipertahankan dalam dimensi penuh (24 atribut), terdapat banyak variabel biologis yang saling bertumpukan (*overlapping information*) dan memiliki ketergantungan multilinear yang kuat, yang justru dapat mengaburkan batas pemisahan kelas pada setiap percabangan pohon keputusan.

Sebagai contoh konkret, parameter kadar urea darah (*blood urea*) dan kreatinin serum (*serum creatinine*) secara medis merupakan produk sisa metabolisme otot yang sama-sama disaring oleh glomerulus ginjal. Memasukkan kedua nilai ini secara mentah dan bersamaan ke dalam model klasifikasi tanpa proses eliminasi akan memicu fenomena perhitungan ganda (*double counting*) informasi yang memiliki korelasi tinggi, sehingga mengakibatkan pembengkakan varians dan penurunan efisiensi pohon keputusan. Lewat mekanisme reduksi fitur teroptimasi hingga menyisakan 12 atribut klinis esensial, mata rantai ketergantungan multilinear tersebut berhasil diputus. Proses ini mengeliminasi variabel redundan tanpa kehilangan esensi informasi prediktifnya, sehingga setiap estimator pohon dalam *Random Forest* dapat bekerja secara maksimal dan objektif pada subset data yang bersih. Lebih jauh lagi, penemuan dalam penelitian ini menyoroti signifikansi klinis yang tinggi terkait daftar subset fitur yang secara konsisten dipertahankan seperti *Specific Gravity* (sg), *Albumin* (al), dan *Hemoglobin* (hemo). Ditinjau dari sudut pandang patofisiologi kedokteran, pemilihan parameter ini sangatlah valid dan rasional. Penurunan nilai *Specific Gravity* secara persisten mengindikasikan hilangnya kemampuan tubulus ginjal dalam melakukan konsentrasi urin. Sementara itu, kemunculan *Albumin* di dalam urin (*albuminuria*) merupakan indikator utama adanya kerusakan struktural pada dinding membran filtrasi glomerulus. Di sisi lain, penurunan drastis kadar *Hemoglobin* mencerminkan gangguan sekresi hormon eritropoetin oleh jaringan ginjal yang rusak, yang memicu kondisi anemia kronis pada pasien.

Dengan memfokuskan proses komputasi hanya pada parameter-parameter kunci ini, model berhasil mengeliminasi atribut pengganggu yang kurang diskriminatif seperti usia (*age*) dan riwayat bakteri (*ba*). Pemurnian fitur serta implementasi transformasi skala berbasis *Z-score* secara langsung menghindarkan model dari jebakan fenomena *overfitting*, yaitu suatu kondisi patologis komputasi di mana model terlalu detail mempelajari karakteristik *noise* yang melekat pada data latih, sehingga kehilangan kemampuan generalisasi ketika dihadapkan pada data uji independen yang baru. Aspek kekuatan lain dari model yang dikembangkan dalam penelitian ini terletak pada tingginya tingkat stabilitas dan konsistensi performa di sepanjang pengujian pada skema 10-Fold Cross Validation. Visualisasi Kurva ROC yang melonjak tajam membentuk sudut siku-siku ideal di pojok kiri atas membuktikan keandalan grafik diskriminasi model secara riil dengan nilai AUC mencapai 0.999 untuk *Random Forest* dan 0.994 untuk *kNN*. Keberhasilan mencatatkan nilai capaian akhir Akurasi global (CA) sebesar 99,00% serta nilai *Recall* mencapai 98,80% pada *Random Forest*, yang didukung oleh temuan *Confusion Matrix* berupa angka eror *False Negative* sebesar 0.0% (250 sampel pasien *ckd* terklasifikasi sempurna), memberikan justifikasi kuat bagi model ini untuk direkomendasikan sebagai fondasi komputasi utama dalam pengembangan aplikasi *Clinical Decision Support System* (CDSS). Dalam skenario klinis medis riil, orientasi utama dari sistem pakar diagnostik adalah memaksimalkan sensitivitas atau nilai *Recall* guna meminimalisir nilai *False Negative*. Kesalahan dalam mendiagnosis pasien sehat sebagai sakit (*False Positive*) hanya akan memicu pemeriksaan laboratorium ulang (*re-testing*) yang berakibat pada pembengkakan biaya minor. Sebaliknya, kesalahan dalam mengklasifikasikan pasien yang sebenarnya menderita ginjal kronis sebagai individu sehat (*False Negative*) adalah kecacatan diagnosis yang fatal karena merampas kesempatan pasien untuk mendapatkan intervensi medis dini sebelum jaringan ginjal mengalami kerusakan permanen (*End-Stage Renal Disease*).

Jika disandingkan dengan studi literatur terdahulu, hasil penelitian ini memberikan kontribusi keilmuan yang signifikan dalam memperluas khazanah optimasi klasifikasi medis domestik. Model *Proposed Random Forest* (99,00%) terbukti mampu melampaui performa model optimasi parameter milik Khasan Galuh Ramadhan dkk [4]. (*Random Forest* + *GridSearchCV*) yang berada di angka 98,75%. Sementara itu, lompatan performa radikal juga ditunjukkan oleh algoritma *kNN* yang melesat hingga 96,00%, mengungguli K-NN standar milik Miftahul Rizal dkk [6]. (91,50%), JST *Backpropagation* milik Chikita Indah Felisa dkk [8]. (92,00%), serta Hibrida SVM-KNN milik Alvin Tarisa Akbar dkk [7]. (94,25%). Keterbatasan fundamental dari algoritma-algoritma klasik tersebut berhasil dijawab secara tuntas dalam penelitian ini melalui skema reduksi 12 fitur klinis esensial. Temuan empiris ini sekaligus menegaskan sebuah paradigma baru dalam *data science* medis, bahwa tahapan seleksi dan standardisasi fitur tidak boleh lagi dipandang remeh sebagai proses opsional tambahan, melainkan harus diposisikan sebagai fase mutlak yang krusial demi membangun arsitektur kecerdasan buatan medis yang andal, akurat, presisi, dan aman bagi keselamatan jiwa pasien.

4. KESIMPULAN

Berdasarkan hasil eksperimen dan analisis komparatif yang telah dilakukan, dapat disimpulkan bahwa arsitektur klasifikasi penyakit ginjal kronis berbasis model *Random Forest* dan K-Nearest Neighbor (kNN) terbukti memberikan performa diagnosis yang sangat superior. Implementasi alur penanganan data (*data pipeline*) melalui standarisasi skala berbasis *Z-score* dan reduksi dimensi dari 24 menjadi 12 atribut klinis esensial terbukti efektif mengeliminasi fenomena perhitungan ganda (*double counting*) akibat ketergantungan multilinear variabel prediktor. Hal ini didukung secara kuat oleh pembuktian patofisiologi kedokteran terhadap pemilihan parameter kunci seperti *Specific Gravity* (sg), *Albumin* (al), dan *Hemoglobin* (hemo) yang secara konsisten mampu menjaga integritas informasi klinis tanpa memicu kendala *overfitting*. Pengujian performa global menggunakan skema 10-Fold Cross Validation menunjukkan dominasi mutlak dari algoritma *Proposed Random Forest* yang berhasil mencatatkan nilai akurasi tertinggi sebesar 99,00%. Validitas performa diskriminasi model ini juga dipertegas oleh capaian nilai AUC sebesar 0.999 serta tingkat sensitivitas klinis yang luar biasa melalui nilai eror *False Negative* sebesar 0.0% pada evaluasi *Confusion Matrix*. Di sisi lain, model *baseline* kNN juga mengalami lompatan performa radikal dengan raihan akurasi mencapai 96,00% dan AUC 0.994. Konfrontasi data lintas literatur secara empiris membuktikan bahwa kedua arsitektur yang diajukan dalam penelitian ini sukses memecahkan rekor performa domestik dengan mengungguli metode-metode mutakhir terdahulu, termasuk *Random Forest* teroptimasi GridSearchCV (98,75%), Hibrida SVM-KNN (94,25%), dan JST *Backpropagation* (92,00%). Dengan demikian, integrasi model ini sangat valid dan direkomendasikan sebagai fondasi utama *Clinical Decision Support System* (CDSS) demi menjamin keselamatan jiwa pasien melalui intervensi medis dini yang aman, akurat, dan presisi.

UCAPAN TERIMAKASIH

Penulis menyampaikan rasa terima kasih dan apresiasi yang setinggi-tingginya kepada seluruh pihak yang telah memberikan dukungan, bantuan, serta motivasi moral maupun material demi kelancaran dan terlaksananya penelitian komparatif ini dengan baik. Ucapan terima kasih secara khusus ditujukan kepada segenap civitas akademika **Universitas HKBP Nommensen P. Siantar** atas segala fasilitas akademis dan iklim riset yang kondusif yang telah diberikan kepada penulis selama masa studi. Tidak lupa, apresiasi yang tulus juga penulis sampaikan kepada keluarga tercinta serta rekan-rekan mahasiswa di lingkungan kampus yang senantiasa memberikan semangat, inspirasi, dan ruang diskusi yang sangat suportif, sehingga seluruh tahapan eksperimen penambangan data ini dapat diselesaikan dengan baik dan tepat waktu.

REFERENCES

- [1] A. Francis *et al.*, "Chronic kidney disease and the global public health agenda: an international consensus," *Nat. Rev. Nephrol.*, vol. 20, no. 7, pp. 473–485, 2024, doi: 10.1038/s41581-024-00820-6.
- [2] A. Novandra *et al.*, "Menjalani Hemodialisis Pendahuluan Penyakit merupakan Ginjal Kronis (PGK) global albumin urine per gram dari creatinin urin), Glomerular Filtration Rate (GFR) < 60ml / menit / 1 , 73 m 2 dengan jangka waktu PGK menyebabkan perubahan masalah kesehatan," vol. 10, pp. 235–253, 2025.
- [3] J. Pendidikan, "Pemanfaatan Artificial Intelligence Dalam Bidang Kesehatan Ganis," vol. 11, no. 1, pp. 234–242, 2024.
- [4] Khasan Galuh Ramadhan, Gentur Wahyu Nyipto Wibowo, and Nadia Annisa Maori, "Komparasi Deteksi Penyakit Ginjal Kronis Menggunakan Algoritma Support Vector Machine Dan Random Forest," *Jurnal Informatika dan Rekayasa Elektronik*, vol. 8, no. 1, pp. 13–21, 2025, doi: 10.36595/jire.v8i1.1288.
- [5] S. Dhara, "Data Pre-Processing and Feature Engineering for Optimization," *Artificial Intelligence Techniques in Mathematical Modeling and Optimization*, pp. 32–56, 2026, doi: 10.1201/9781003633402-3.
- [6] M. Rizal, M. Z. Syahaf, S. R. Priyambodo, and Y. Rhamdani, "Optimasi Algoritma Naïve Bayes Menggunakan Forward Selection Untuk Klasifikasi Penyakit Ginjal Kronis," *Naratif: Jurnal Nasional Riset, Aplikasi dan Teknik Informatika*, vol. 5, no. 1, pp. 71–80, 2023, doi: 10.53580/naratif.v5i1.200.
- [7] A. T. Akbar, N. Yudistira, and A. Ridok, "Identifikasi Gagal Ginjal Kronis dengan Mengimplementasikan Metode Support Vector Machine beserta K-Nearest Neighbour (SVM-KNN)," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 10, no. 2, pp. 301–308, 2023, doi: 10.25126/jtiik.20236059.
- [8] D. Irmansyah, "Jurnal JISIILKOM (Jurnal Inovasi Sistem Informasi & Ilmu Komputer) Prediksi Penyakit Paru-Paru Menggunakan Jaringan Syaraf Tiruan Dengan Metode Backpropagation Article Info," *Jurnal JISIILKOM*, vol. 3, no. 1, pp. 3025–4868, 2025.
- [9] G. Anantasia and S. R. Rindrayani, "Metodologi Penelitian Quasi Eksperimen," *Adiba: Journal of Education*, vol. 5, no. 2, pp. 183–192, 2025.
- [10] F. R. Aftha Harianto, Z. Alawi, and I. A. Sa'ida, "Pengaruh Komposisi Split Data Pada Akurasi Klasifikasi Penderita Diabetes Menggunakan Algoritma Machine Learning," *Jurnal Sistem Informasi dan Informatika (Simika)*, vol. 8, no. 1, pp. 36–44, 2025, doi: 10.47080/simika.v8i1.3663.
- [11] P. Aisyiyah and R. Devi, "JIPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika) Journal homepage: <https://jurnal.stkipppgritulungagung.ac.id/index.php/jipi> Klasifikasi Penyakit Gagal Ginjal Kronis Dengan Metode KNN (Studi Kasus RS Di Kab Gresik)," vol. 9, no. 3, pp. 1739–1748, 2024, [Online]. Available: <https://doi.org/10.29100/jipi.v9i3.6226>
- [12] I. Wisnuadji Gamadarenda and I. Waspada, "Implementation of Data Mining for the Detection of Chronic Kidney Disease (Ckd) Using K-Nearest Neighbor (Knn) With Backward Elimination," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, vol. 7, no. 2, pp. 417–426, 2020, doi: 10.25126/jtiik.202071896.
- [13] F. M. Natsir, R. Y. Bakti, and T. Wahyuni, "Analisis Deteksi Dini Penyakit Jantung dengan Pendekatan Support Vector Machine pada Data Pasien," *Arus Jurnal Sains dan Teknologi*, vol. 2, no. 2, pp. 437–446, 2024, doi: 10.57250/ajst.v2i2.669.

- [14] L. Sari, A. Romadloni, and R. Listyaningrum, "Penerapan Data Mining dalam Analisis Prediksi Kanker Paru Menggunakan Algoritma Random Forest," *Infotekmesin*, vol. 14, no. 1, pp. 155–162, 2023, doi: 10.35970/infotekmesin.v14i1.1751.
- [15] N. P. Nugraha, R. Azim, S. Z. Daffa, and P. S. Ningayu, "Perbandingan Akurasi Metode Naïve 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.
- [16] A. Ar, R. Hasibuan, and N. A. Hrp, "Implementasi Algoritma Decision Tree untuk Klasifikasi Penyakit Ginjal Kronis Berbasis Machine Learning," vol. 4, no. 2, pp. 49–61, 2026.
- [17] F. Hafizulhaq and A. Andasuryani, "Orange Classification using Naïve Bayes and K-Nearest Neighbor Algorithms based on Its Physical Properties," *EKSAKTA: Journal of Sciences and Data Analysis*, vol. 7, no. 1, pp. 197–205, 2026, [Online]. Available: <https://journal.uui.ac.id/Eksakta/article/view/43102>
- [18] O.- Pahlevi, A.- Amrin, and Y.- Handrianto, "Implementasi Algoritma Klasifikasi Random Forest Untuk Penilaian Kelayakan Kredit," *Jurnal Infortech*, vol. 5, no. 1, pp. 71–76, 2023, doi: 10.31294/infortech.v5i1.15829.
- [19] Z. Z. Hulaifah Al Abrori and E. R. Subhiyakto, "Analisis Komparatif Akurasi Prediksi Kanker Payudara Menggunakan Algoritma Random Forest dan Logistic Regression," *Jurnal Algoritma*, vol. 22, no. 1, pp. 300–311, 2025, doi: 10.33364/algoritma/v.22-1.2164.
- [20] F. Putra, H. F. Tahiyat, R. M. Ihsan, R. Rahmadden, and L. Efrizoni, "Penerapan Algoritma K-Nearest Neighbor Menggunakan Wrapper Sebagai Preprocessing untuk Penentuan Keterangan Berat Badan Manusia," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 1, pp. 273–281, 2024, doi: 10.57152/malcom.v4i1.1085.