

Prediction of Chronic Kidney Disease Using the XGBoost Classifier Algorithm

[Prediksi Penyakit Ginjal Kronik Menggunakan Algoritma Xgboost Classifier]

Bayu Dimas Abimanyu¹⁾, Hamzah Setiawan, S.Kom., M.Kom²⁾, Ade Eviyanti, S.Kom., M.Kom³⁾, Ika Ratna Indra Astutik, S.Kom., MT⁴⁾.

¹⁾Program Studi Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

²⁾ Program Studi Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

*Email Penulis Korespondensi: bayudimasa1999@gmail.com, hamzah@umsida.ac.id

Abstract. *Chronic kidney disease is a medical condition marked by a slow, ongoing, and progressive deterioration of kidney function. When it is not identified early, the disease may bring about serious complications that greatly reduce a patient's quality of life, and it can ultimately progress to end-stage renal failure that necessitates dialysis or kidney transplantation. For that reason, establishing a reliable early-detection system is crucial to enable timely diagnosis and treatment. This study sets out to build a chronic kidney disease prediction model using the XGBoost Classifier algorithm, which is widely recognized for strong performance in classification problems and for its capacity to work with complex datasets. The research is carried out through multiple steps: preprocessing the data to clean and ready the dataset, dividing it into training and testing subsets, training the model, and evaluating its performance. Model quality is assessed with a confusion matrix and a classification report, which together quantify accuracy, precision, recall, and the F1-score. The experimental outcomes indicate that the proposed model reaches an accuracy of 98.78%, reflecting outstanding predictive capability. In addition, the precision, recall, and F1-score metrics for both classes remain consistently high. The confusion matrix reports 53 True Negatives, 28 True Positives, 1 False Positive, and 0 False Negatives, pointing to a very low error rate. Overall, these results imply that the XGBoost-based model is highly effective at classifying chronic kidney disease cases and holds strong promise for use as a decision support tool for early detection within healthcare systems.*

Keywords – Chronic Kidney Disease, XGBoost Classifier, Machine Learning, Classification, Disease Prediction.

Abstrak Penyakit ginjal kronik adalah kondisi medis yang berlangsung bertahap dan semakin progresif, ditandai oleh menurunnya fungsi ginjal dalam jangka waktu panjang. Jika tidak dikenali sejak dini, kondisi ini dapat memicu komplikasi berat yang merusak kualitas hidup pasien, bahkan berpotensi berkembang hingga gagal ginjal tahap akhir yang menuntut terapi cuci darah atau transplantasi ginjal. Karena itu, dibutuhkan sebuah sistem prediksi yang dapat mempercepat sekaligus meningkatkan ketepatan proses deteksi dini. Penelitian ini disusun untuk membangun model prediksi penyakit ginjal kronik dengan algoritma XGBoost Classifier, yang dikenal unggul dalam menangani data klasifikasi. Alur penelitian mencakup tahapan preprocessing untuk membersihkan dan menyiapkan dataset, pembagian dataset menjadi data training dan testing, pelatihan model, serta penilaian kinerja melalui confusion matrix dan classification report. Berdasarkan hasil pengujian, model yang dikembangkan mampu mencapai akurasi 98,78%. Nilai precision, recall, dan F1-score yang diperoleh juga memperlihatkan performa sangat baik pada kedua kelas, baik kelas positif maupun kelas negatif. Confusion matrix yang dihasilkan mencatat 53 True Negative, 28 True Positive, 1 False Positive, dan 0 False Negative. Temuan ini menunjukkan bahwa model memiliki tingkat kesalahan yang sangat rendah serta kemampuan yang kuat untuk mengklasifikasikan data pasien secara tepat. Dengan demikian, model XGBoost Classifier ini berpotensi dimanfaatkan sebagai alat bantu dalam sistem pendukung keputusan untuk deteksi dini penyakit ginjal kronik.

Kata Kunci - Penyakit Ginjal Kronik, XGBoost Classifier, Machine Learning, Klasifikasi, Prediksi Penyakit.

I. PENDAHULUAN

Penyakit ginjal kronik (PGK) kini menjadi salah satu isu kesehatan yang kian mengganggu perhatian dunia. PGK merupakan kondisi medis yang ditandai dengan penurunan fungsi ginjal yang berlangsung perlahan namun pasti, sering kali tanpa gejala yang nyata pada fase awal. Namun, ketika penyakit ini telah masuk ke tahap yang lebih lanjut, PGK dapat memicu berbagai komplikasi serius, termasuk gagal ginjal yang membutuhkan penanganan melalui dialisis atau transplantasi ginjal [1].

Tingkat kejadian PGK semakin meningkat, dan faktor-faktor risiko yang berkaitan dengan pola hidup modern telah diidentifikasi sebagai penyebab utama [2]. Gaya hidup masyarakat modern yang kerap tidak sehat—misalnya pola makan dengan kandungan lemak dan garam yang tinggi, kurangnya aktivitas fisik yang memadai, serta kebiasaan mengonsumsi alkohol dan merokok, ditambah minimnya asupan air putih dapat meningkatkan risiko penyakit ginjal kronik (PGK) [3]. Di sisi lain, penggunaan obat-obatan tertentu, khususnya obat antiinflamasi nonsteroid (OAINS) yang lazim dipakai untuk meredakan nyeri dan peradangan, berpotensi menimbulkan kerusakan ginjal bila digunakan dalam jangka panjang serta pada dosis tinggi [4]. Selain itu, minimnya pemantauan kesehatan secara rutin oleh masyarakat membuat kondisi kesehatan, termasuk PGK, sering kali baru terdeteksi ketika sudah terlambat, sehingga upaya pencegahan dan intervensi yang efektif menjadi terhambat [5].

Di Indonesia, penyakit ginjal kronik (PGK) juga menjadi persoalan kesehatan yang patut mendapat perhatian sungguh-sungguh. Mengacu pada temuan Riset Kesehatan Dasar (Riskesdas) 2018 yang disusun oleh Badan Penelitian dan Pengembangan Kesehatan Kementerian Kesehatan Republik Indonesia, prevalensi PGK di Indonesia tercatat sebesar 0,38%—kurang lebih setara dengan 3,8 kasus per 1.000 penduduk. Dari total tersebut, sekitar 60% penderita PGK stadium lanjut memerlukan terapi dialisis. Besaran prevalensi ini tidak sama di setiap provinsi: Provinsi Kalimantan Utara mencatat angka tertinggi sebesar 0,64%, sedangkan Provinsi Sulawesi Barat menjadi yang terendah dengan 0,18%. Walau angka Riskesdas 2018 tergolong lebih rendah dibandingkan beberapa negara lain, penelitian Perhimpunan Nefrologi Indonesia (PERNEFRI) pada tahun 2006 justru memperlihatkan prevalensi yang lebih tinggi, yakni 12,5%. Sampai saat ini, ketersediaan data insidensi dan prevalensi PGK pada anak di Indonesia masih terbatas, tetapi pada tahun 2017 tercatat 220 anak dengan PGK tahap akhir yang menjalani dialisis dan 13 anak yang menjalani transplantasi ginjal di 16 rumah sakit pendidikan di Indonesia. Temuan-temuan tersebut menegaskan bahwa PGK adalah masalah kesehatan yang bermakna dan menuntut upaya deteksi serta penanganan yang lebih efektif [6].

Riset tentang penerapan machine learning untuk klasifikasi telah memperlihatkan capaian yang sangat memuaskan pada beragam situasi. Setiawan, Fatichah, dan Saikhu (2023) merancang sebuah sistem klasifikasi multilabel berbasis data umpan balik mahasiswa dengan memanfaatkan representasi teks Bidirectional Encoder Representations from Transformers (BERT) yang dipadukan dengan beberapa algoritma machine learning, yakni Support Vector Machine (SVM), K-Nearest Neighbors (KNN), Random Forest (RF), dan Decision Tree (DT). Temuan penelitian ini mengungkapkan bahwa metode SVM berkernel linear menjadi yang paling unggul, dengan akurasi 82% dan F1-score 90%, sehingga menegaskan bahwa algoritma machine learning dapat menghasilkan performa klasifikasi yang baik bahkan untuk data yang kompleks [7].

Penelitian lain yang dilakukan oleh Ade Eviyanti, Saikhu, dan Fatichah (2022) membandingkan metode machine learning dan deep learning untuk mendeteksi kejang epilepsi menggunakan sinyal Electroencephalogram (EEG). Penelitian itu menggunakan ekstraksi fitur melalui Discrete Wavelet Transform, lalu menguji dan membandingkan beberapa algoritma, yakni SVM, KNN, Random Forest, Decision Tree, serta Long Short-Term Memory. Dari rangkaian pengujian yang dilakukan, SVM tampil sebagai yang paling unggul, karena akurasi, precision, recall, dan F1-score masing-masing berhasil mencapai 100%. Temuan ini menegaskan bahwa algoritma pembelajaran mesin mempunyai kapasitas yang sangat tinggi dalam melakukan klasifikasi pada data Kesehatan [8].

Penerapan kecerdasan buatan juga telah dimanfaatkan pada sistem pendukung keputusan di bidang kesehatan. Ika Ratna Indra Astutik, Indahyanti, dan Rinata (2023) mengembangkan aplikasi deteksi dini stunting berbasis Forward Chaining Inference Machine untuk membantu masyarakat mengenali risiko stunting sejak dini. Hasil penelitian menunjukkan bahwa aplikasi tersebut mampu membantu orang tua maupun tenaga kesehatan dalam melakukan deteksi awal sehingga tindakan pencegahan dapat dilakukan lebih cepat [9].

Khusus pada penyakit ginjal kronik, penelitian terbaru oleh nimithap dkk. (2026) mengembangkan model prediksi perkembangan Chronic Kidney Disease (CKD) menggunakan algoritma Random Forest pada data rumah sakit. Penelitian tersebut memperoleh akurasi sebesar 96,25% dengan ROC-AUC sebesar 1,00 serta menunjukkan bahwa hemoglobin, serum kreatinin, packed cell volume, dan specific gravity merupakan faktor yang paling berpengaruh terhadap prediksi CKD. Penelitian tersebut juga memanfaatkan SHAP (SHapley Additive exPlanations), sehingga model yang dihasilkan menjadi lebih mudah dipahami dan dapat mendukung pengambilan keputusan klinis [10].

Walaupun sejumlah penelitian terdahulu telah membuktikan bahwa algoritma machine learning dapat mencapai performa klasifikasi yang tinggi pada beragam permasalahan kesehatan, kajian yang secara khusus meneliti prediksi penyakit ginjal kronik dengan algoritma XGBoost Classifier masih tergolong sedikit. Banyak penelitian CKD sebelumnya lebih sering memakai algoritma seperti Naive Bayes, Support Vector Machine, K-Nearest Neighbor, maupun Random Forest. Padahal, XGBoost unggul berkat kemampuannya mengolah data yang kompleks, penerapan teknik regularisasi untuk menekan terjadinya overfitting, penanganan missing values, serta pelatihan yang lebih efisien melalui parallel processing. Karena itu, penelitian ini menggunakan algoritma XGBoost Classifier untuk membangun model prediksi penyakit ginjal kronik yang ditargetkan mampu menghasilkan tingkat akurasi yang tinggi.

Penelitian ini disusun untuk merancang model prediksi yang dapat mengenali risiko penyakit PGK pada individu berdasarkan berbagai faktor yang memengaruhinya, termasuk kebiasaan masyarakat. Adapun salah satu algoritma pembelajaran mesin yang akan diterapkan dalam penelitian ini adalah algoritma XGBoost Classifier [11]. Algoritma ini memiliki sejumlah kelebihan yang signifikan [12]. Pertama, XGBoost dikenal dengan tingkat akurasi prediksi yang tinggi, sebuah aspek kritis dalam penelitian kesehatan seperti deteksi dini penyakit ginjal kronik (PGK). Kemampuan algoritma ini untuk menghasilkan prediksi yang akurat akan sangat berarti dalam mengidentifikasi individu yang berisiko PGK lebih awal [13].

Algoritma XGBoost, yang merupakan kepanjangan dari Extreme Gradient Boosting, bekerja dengan cara membangun serangkaian pohon keputusan secara berurutan guna meningkatkan tingkat keakuratannya. Langkah ini bermula ketika pohon keputusan pertama dibuat, lalu pohon tersebut menelusuri dan mengenali pola-pola yang terdapat di dalam data [13]. Selanjutnya, pohon-pohon berikutnya dibangun dengan fokus pada koreksi kesalahan prediksi model sebelumnya. Gradien dari kesalahan prediksi ini digunakan untuk memberikan bobot pada setiap instance dalam data, memberikan penekanan lebih pada kasus-kasus dengan kesalahan prediksi tinggi.

Algoritma ini menggunakan teknik regularisasi untuk mencegah overfitting, mengoptimalkan kompleksitas model, dan memastikan generalisasi yang baik pada data baru. Hasil prediksi dari setiap pohon digabungkan dengan memberikan bobot pada masing-masing prediksi, menghasilkan prediksi akhir. Teknik optimasi gradien turun digunakan untuk menemukan bobot optimal dengan menyesuaikannya pada setiap iterasi berdasarkan gradien dari fungsi kesalahan.

Selain itu, XGBoost dapat mengelola nilai yang hilang (missing values) dalam data, sekaligus menyediakan mekanisme early stopping untuk menahan overfitting dengan menghentikan pelatihan saat kinerja model tidak menunjukkan peningkatan yang berarti. Algoritma ini juga memungkinkan pemrosesan paralel, sehingga proses pelatihan dapat berlangsung lebih cepat—terutama ketika berhadapan dengan dataset berukuran besar. Berkat sejumlah keunggulan tersebut, XGBoost menjadi salah satu algoritma yang efektif untuk menangani tugas klasifikasi, termasuk pada penelitian ini yang bertujuan memprediksi risiko Penyakit Ginjal Kronik (PGK).

Penelitian ini memanfaatkan data publik yang bersumber dari UCI Machine Learning Repository, dengan jumlah total 400 data pasien. Di dalam himpunan data, terdapat dua kategori, yakni CKD (penyakit ginjal kronik) dan Non-CKD. Kategori CKD menggambarkan pasien yang terindikasi mengalami penyakit ginjal kronik, yaitu penurunan fungsi ginjal yang terjadi secara progresif dan berlangsung dalam jangka panjang. Adapun kategori Non-CKD merujuk pada pasien yang tidak terindikasi menderita penyakit ginjal kronik. Mengacu pada komposisi distribusi data, sebanyak 62,5% data termasuk dalam kelas CKD, sementara 37,5% data masuk ke kelas Non-CKD.

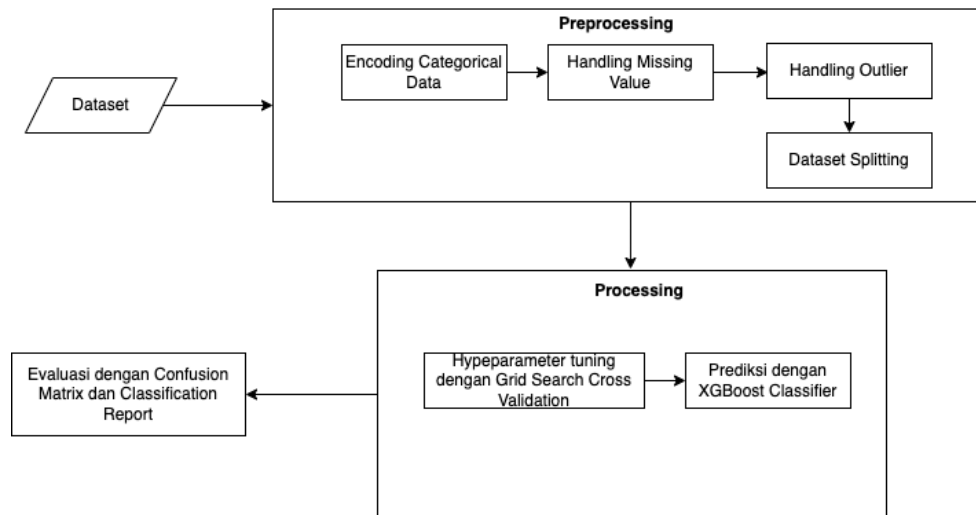
Dengan mengoptimalkan keunggulan-keunggulan tersebut, penelitian ini diarahkan untuk merancang alat prediksi yang presisi dan dapat dipercaya guna mendeteksi risiko penyakit ginjal kronik lebih dini berdasarkan faktor-faktor risiko yang relevan. Penelitian ini juga diharapkan mampu menyumbang secara berarti bagi upaya pencegahan serta pengelolaan penyakit ginjal kronik, sehingga pada akhirnya dapat meningkatkan kualitas hidup dan kesejahteraan masyarakat secara menyeluruh.

Berdasarkan uraian latar belakang di atas, peneliti terdorong untuk membahas sekaligus menelaah secara lebih mendalam penelitian yang membahas penyakit ginjal kronik dengan judul “Prediksi Penyakit Ginjal Kronik Menggunakan Algoritma XGBoost Classifier”.

II. METODE

1. Tahapan Penelitian

Subbab ini menjelaskan kerangka umum tahapan penelitian yang digunakan dalam proses pengembangan model prediksi penyakit ginjal kronik. Penelitian ini dirancang dengan alur yang terstruktur, dimulai dari pengolahan data lalu berlanjut pada evaluasi kinerja model, sebagaimana diperlihatkan pada Gambar 1



Gambar 1 Tahapan Penelitian

Mengacu pada Gambar 1, rangkaian penelitian dimulai dari pemrosesan data mentah, yakni pengelolaan data awal yang berhasil dihimpun dari sumber dataset. Berikutnya dilakukan pra-pemrosesan data yang mencakup pengodean data kategoris, penanganan nilai yang hilang (missing values), serta penertiban nilai ekstrem (outlier) demi menajamkan kualitas data. Sesudah itu, dataset dipisahkan menjadi data pelatihan (training data) dan data pengujian (testing data) untuk sekaligus membangun serta mengevaluasi model.

Tahap berikutnya adalah pemodelan, yang mencakup penyetelan hyperparameter menggunakan metode GridSearch dengan teknik cross-validation, serta implementasi algoritma XGBoost Classifier untuk melakukan prediksi. Tahap penutup berupa evaluasi model dengan memanfaatkan confusion matrix dan classification report untuk menilai performanya melalui metrik akurasi, presisi, recall, serta F1-score. Setiap langkah dalam proses ini sama pentingnya, karena bersama-sama mereka membentuk model prediktif yang akurat dan dapat dipercaya.

2. Data

Pada subbab ini dipaparkan data yang dijadikan masukan dalam perangkat lunak penelitian. Data tersebut selanjutnya akan diolah dan diuji sehingga menghasilkan keluaran yang selaras dengan sasaran penelitian. Sumber data yang digunakan berasal dari UCI Machine Learning Repository, dengan total 400 data pasien. Kumpulan data ini dibagi menjadi dua kelas, yaitu CKD dan Not-CKD, dengan komposisi 62,5% data berada pada kelas CKD dan 37,5% data berada pada kelas Not-CKD. Adapun rincian atribut yang dipakai disajikan pada tabel berikut.

Tabel 1 Tabel Atribut

Atribut	Keterangan
age	Usia pasien dengan rentang 2–90 tahun
bp	Tekanan darah (mmHg) dengan nilai 50–180
Sg	Berat jenis spesifik urin dengan nilai 1.005–1.025
Al	Kadar albumin urin dengan nilai 0–5

Su	Kadar gula urin dengan nilai 0–5
rbc	Kondisi sel darah merah (normal, abnormal)
Pc	Kondisi sel leukosit (normal, abnormal)
pcc	Gumpalan sel leukosit (present, notpresent)
Ba	Keberadaan bakteri (present, notpresent)
bgr	Glukosa darah acak (mg/dL) dengan rentang 22–490
bu	Urea darah (mg/dL) dengan rentang 1.5–391
Sc	Kreatinin serum (mg/dL) dengan rentang 0.4–76.0
sod	Kadar natrium (mEq/L) dengan rentang 104–163
pot	Kadar kalium (mEq/L) dengan rentang 2.5–7.6
hemo	Kadar hemoglobin (g/dL) dengan rentang 3.1–17.8
pcv	Volume sel darah merah (%) dengan rentang 9–54
wc	Jumlah sel darah putih (/mm ³) dengan rentang 2.200–26.400
Rc	Jumlah sel darah merah (juta sel/mm ³) dengan rentang 2.1–8.0
htn	Riwayat hipertensi (yes, no)
dm	Riwayat diabetes mellitus (yes, no)
cad	Riwayat penyakit arteri koroner (yes, no)
appet	Nafsu makan (good, poor)
Pe	Pembengkakan kaki (yes, no)
ane	Kondisi anemia (yes, no)
classification	Kelas data (ckd, notckd)

III. HASIL DAN PEMBAHASAN

Pada tahap pra-pemrosesan data, para peneliti menggunakan teknik Label Encoder untuk menghadapi variabel kategoris yang terdapat di dalam dataset. Cara ini mempermudah proses dengan mengubah data kategoris menjadi bentuk yang dapat diolah oleh algoritma machine learning.

1. *Preprocessing*

Preprocessing, atau pra-pemrosesan data, adalah tahap awal penting dalam analisis data dan pembuatan model machine learning [14]. Tahap ini melibatkan penyiapan dan transformasi data mentah untuk memastikan kualitas dan kesesuaian data sebelum dilakukan analisis lebih lanjut. Ini termasuk langkah-langkah penting seperti pengkodean data kategori, penanganan data yang hilang, dan pengelolaan data ekstrem (outlier). Preprocessing yang efektif sangat memengaruhi hasil penelitian berbasis data dan pembuatan model prediktif yang akurat.

a. *Encoding Categorical Data*

Pada tahap pra-pemrosesan data, peneliti menggunakan teknik Label Encoder untuk menghadapi variabel kategoris yang terdapat di dalam dataset. Dengan strategi ini, data kategoris diarahkan agar berubah menjadi bentuk yang siap diolah oleh algoritma machine learning [15]. Dengan Label Encoder, setiap nilai kategori pada kolom yang relevan diubah menjadi label numerik yang unik. Transformasi ini membuat data lebih selaras untuk proses pemodelan, sehingga peneliti dapat memanfaatkan informasi kategoris dalam model tanpa harus membuat banyak kolom tambahan yang berpotensi mengubah dimensi kumpulan data secara besar. Karena itu, teknik ini menjadi tahapan penting dalam persiapan data, yang membantu peneliti meningkatkan akurasi serta performa model mereka.

Sebagai contoh, pada dataset Penyakit Ginjal Kronik terdapat variabel kategoris *pc* (sel leukosit) yang memiliki nilai *normal* dan *abnormal*. Karena algoritma machine learning hanya mampu mengolah data berbentuk angka, nilai kategoris tersebut perlu diubah menjadi bilangan melalui Label Encoder. Misalnya, Label Encoder memberikan kode *normal* = 0 dan *abnormal* = 1. Contoh data sebelum encoding:

Tabel 2 Contoh data sebelum encoding

No	pc
1	normal
2	abnormal
3	normal
4	abnormal
5	normal

Proses encoding dilakukan dengan mengganti setiap nilai kategoris dengan kode numerik sesuai label yang telah ditentukan. Berikut contoh perhitungan encoding:

1. Data ke-1: *pc* = *normal* → 0
2. Data ke-2: *pc* = *abnormal* → 1
3. Data ke-3: *pc* = *normal* → 0
4. Data ke-4: *pc* = *abnormal* → 1
5. Data ke-5: *pc* = *normal* → 0

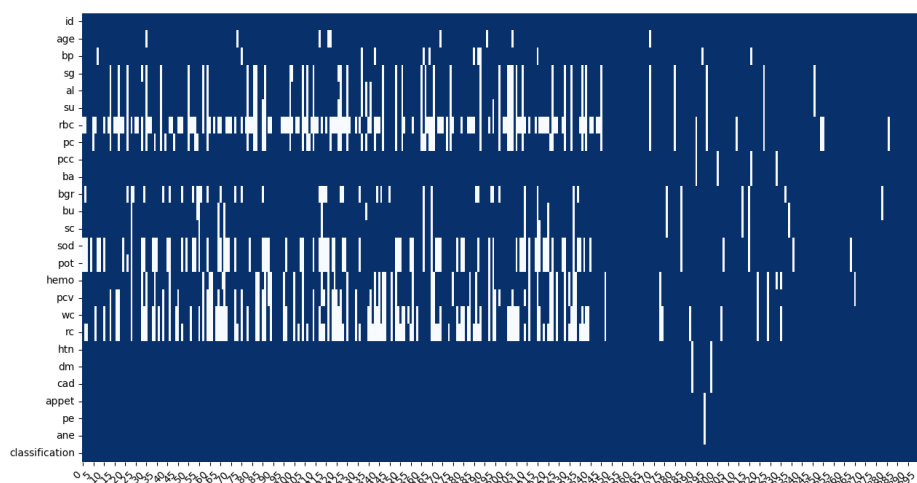
Sehingga hasil setelah encoding adalah:

Tabel 3 Data Hasil Encoding

No	pc (sebelum)	pc (sesudah encoding)
1	normal	0
2	abnormal	1
3	normal	0
4	abnormal	1
5	normal	0

b. Handling Missing Value

Pada fase ini, penelitian berfokus pada upaya menanggulangi nilai-nilai yang hilang di dalam kumpulan data yang dipakai. Visualisasi pada Gambar 3.2 menampilkan persebaran nilai-nilai yang tidak ada dalam dataset yang menjadi pusat perhatian penelitian ini.



Gambar 2 Visualisasi Missing Value

Prinsip-prinsip dalam penanganan data yang hilang sangat krusial dalam memastikan analisis data tetap akurat dan dapat diandalkan. Untuk menanggapi kendala tersebut, penelitian ini menggunakan metode imputasi mean. Metode ini diawali dengan menghitung nilai rata-rata pada tiap kolom yang memiliki data hilang, kemudian mengganti bagian yang kosong dengan rata-rata yang diperoleh. Langkah ini bertujuan melindungi integritas kumpulan data sekaligus menekan dampak yang mungkin ditimbulkan oleh nilai-nilai yang absen terhadap analisis. Melalui imputasi mean yang diterapkan dalam penelitian ini, peneliti berupaya memastikan bahwa data yang dipakai untuk proses pemodelan tetap lengkap, konsisten, dan siap digunakan pada tahap berikutnya.

c. Handling Outlier

Dalam proses penanganan *outlier*, metode yang digunakan adalah Z-Score. Metode Z-Score dipakai untuk menandai data yang menembus batas ketentuan, yaitu di luar rentang -3 sampai +3. Nilai Z-Score menyatakan sejauh mana sebuah angka bergeser dari rata-rata, diukur dengan satuan standar deviasi. Makin tinggi nilai Z-Score, makin jauh pula data tersebut dari rata-rata—sehingga peluangnya untuk berubah menjadi outlier semakin besar. Rumus Z-Score tersaji pada Persamaan (3.1):

$$z = \frac{x - \mu}{\sigma}$$

Keterangan:

Z = nilai Z-Score

X = nilai data

μ = nilai rata-rata (mean)

σ = standar deviasi

Data dikategorikan sebagai *outlier* apabila memiliki nilai $Z > 3$ atau $Z < -3$. Sebagai contoh, digunakan atribut hemo (hemoglobin) dengan data berikut:

8.0, 9.5, 10.2, 11.0, 12.5, 13.0, 15.0, 17.5

Langkah pertama adalah menghitung rata-rata (mean):

$$\mu = \frac{8.0 + 9.5 + 10.2 + 11.0 + 12.5 + 13.0 + 15.0 + 17.5}{8}$$

Misalnya standar deviasi (σ) = 3.05

Kemudian hitung Z-Score untuk nilai 17.5:

$$z = \frac{17.5 - 12.09}{3.05}$$

$$z = \frac{5.41}{3.05}$$

$$z = 1.77$$

Karena nilai Z-Score = 1.77 masih berada dalam rentang -3 hingga +3, maka data tersebut bukan outlier.

d. Dataset Splitting

Pada fase ini, dataset dalam penelitian ini dibagi menjadi dua subset utama, yaitu data pelatihan dan data uji [16]. Proses pembagian data—atau data splitting—adalah cara untuk memecah sebuah dataset menjadi dua bagian atau lebih, sehingga terbentuk subhimpunan data yang terpisah [17]. Umumnya, pembagian tersebut melibatkan dua bagian: bagian pertama dipakai untuk menguji dan mengevaluasi model, sedangkan bagian lainnya digunakan untuk melatih model [18]. Namun, dalam penelitian ini peneliti memutuskan untuk menggunakan tiga variasi pembagian yang berbeda, dan salah satunya adalah:

Tabel 4 Pembagian Rasio Dataset

Rasio Data Tes	Rasio Data Pelatihan
90%	10%
80%	20%
75%	25%

Pengaturan beragam rasio ini dibuat agar model yang tengah dikembangkan memperoleh data yang cukup untuk menghasilkan kualitas pelatihan yang memadai, sekaligus tetap memberi dasar evaluasi yang kokoh terhadap performa model yang akhirnya dihasilkan. Data pelatihan berperan untuk membiasakan model mengenali pola yang tersimpan dalam kumpulan data, sedangkan data uji dipakai untuk menilai seberapa jauh model sanggup memprediksi data baru yang tidak pernah ikut dalam proses pelatihan. Pembagian kumpulan data ini menjadi langkah krusial dalam pengembangan model prediksi penyakit ginjal kronik, serta membantu memastikan penilaian kinerja model dilakukan secara objektif.

2. Processing

a. Klasifikasi dengan XGBoost

Pada fase ini, peneliti memutuskan untuk menggunakan algoritma XGBoost guna merancang model prediksi untuk penyakit ginjal kronik. XGBoost, yang merupakan singkatan dari “Extreme Gradient Boosting,” termasuk salah satu algoritma pembelajaran mesin yang sangat andal, terutama untuk tugas klasifikasi sebagaimana yang menjadi fokus penelitian ini.

Cara kerja XGBoost berawal dari penggabungan banyak pohon keputusan secara adaptif agar terbentuk model yang jauh lebih kuat. Setiap pohon keputusan bekerja sebagai model yang mengambil keputusan berdasarkan serangkaian aturan tertentu. XGBoost lalu melangkah lebih jauh dengan menumbuhkan banyak pohon keputusan secara adaptif, di mana tiap pohon menghasilkan prediksi yang berdiri sendiri. Prediksi dari seluruh pohon kemudian digabungkan untuk menghasilkan prediksi akhir yang lebih akurat [19].

Keunggulan utama XGBoost terletak pada kemampuannya menekan overfitting—yakni pemodelan yang terlalu rinci pada data pelatihan—serta performanya yang luar biasa, khususnya saat berhadapan dengan kumpulan data berukuran besar yang memiliki banyak atribut, sebagaimana diterapkan dalam penelitian ini. Secara matematis, prediksi pada iterasi ke- t dirumuskan sebagai berikut:

$$\hat{y}_i^{(t)} = \hat{y}_i^{(t-1)} + \eta \cdot f_t(x_i)$$

Keterangan:

$\hat{y}_i^{(t)}$ = prediksi pada iterasi ke t

$\hat{y}_i^{(t-1)}$ = prediksi pada iterasi sebelumnya

η = learning rate

$f_t(x_i)$ = fungsi pohon Keputusan pada iterasi ke t

Pada iterasi awal, prediksi biasanya dimulai dari nilai rata-rata target:

$$\hat{y}_0 = \frac{1}{n} \sum_{i=1}^n y_i$$

Kemudian XGBoost menghitung residual atau error menggunakan rumus:

$$r_i = y_i - \hat{y}_i$$

Keterangan:

r_i = residual

y_i = nilai actual

$\hat{y}_i$ = nilai prediksi

Residual ini digunakan untuk membangun pohon keputusan baru yang bertujuan meminimalkan error. Proses ini dilakukan secara berulang hingga model mencapai performa optimal. Misalnya digunakan variabel pc (sel leukosit) dengan encoding:

normal = 0

abnormal = 1

Dataset sederhana:

Tabel 5 Contoh Data

Data	pc	y (aktual)
1	0	0
2	1	1
3	1	1
4	0	0

Langkah 1: Menghitung prediksi awal

Rumus:

$$\hat{y}_0 = \frac{1}{n} \sum_{i=1}^n y_i$$

$$\hat{y}_0 = \frac{0 + 1 + 1 + 0}{4}$$

$$\hat{y}_0 = 0.5$$

Langkah 2: Menghitung residual

Rumus:

$$r_i = y_i - \hat{y}_i$$

Perhitungan:

Tabel 6 Hasil Perhitungan

Data	y	prediksi	residual
1	0	0.5	-0.5
2	1	0.5	0.5
3	1	0.5	0.5
4	0	0.5	-0.5

Langkah 3: Membuat pohon keputusan

Misalnya pohon menghasilkan:

pc abnormal $\rightarrow$ 0.5

pc normal $\rightarrow$ -0.5

Langkah 4: Update prediksi

Rumus:

$$\hat{y}_{baru} = \hat{y}_{lama} + \eta \cdot f(x)$$

Missal learning rate :

$$\eta = 0.5$$

Untuk pc abnormal:

$$\hat{y} = 0.5 + (0.5 \times 0.5)$$

$$\hat{y} = 0.75$$

Untuk pc normal:

$$\hat{y} = 0.5 + (0.5 \times -0.5)$$

$$\hat{y} = -0.25$$

Langkah 5: Menentukan kelas

Rumus fungsi klasifikasi:

$$kelas = \begin{cases} 1, \hat{y} \geq 0.5 \\ 0, \hat{y} < 0.5 \end{cases}$$

Hasil:

Tabel 3. 1 Hasil Prediksi

Data	pc	prediksi	kelas prediksi
1	0	0.25	0
2	1	0.75	1
3	1	0.75	1
4	0	0.25	0

i. Hyperparameter Tuning dengan GridSearchCV

Pada tahap ini, peneliti akan menyetel hyperparameter menggunakan metode Grid Search Cross-Validation (Grid Search CV) untuk memperoleh konfigurasi terbaik dari model XGBoost yang akan dipakai dalam melakukan prediksi penyakit ginjal kronik [20]. Berikut adalah daftar hyperparameter yang akan dioptimalkan.

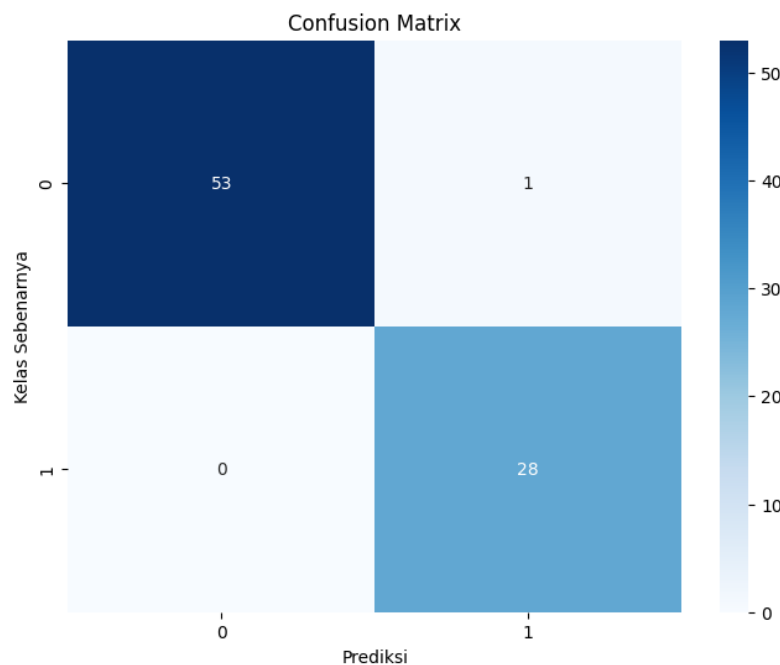
Tabel 3.2 Hyperparameter Tuning

Hyperparameter	Value
n_estimators	100, 200, 300
max_depth	3, 4, 5
Learning_rate	0.01, 0.1, 0.2

Proses penalaan ini diarahkan untuk menemukan kombinasi parameter hiper terbaik guna meningkatkan performa model XGBoost dalam memprediksi penyakit ginjal kronik. Setelah konfigurasi parameter hiper terbaik diperoleh, peneliti kemudian melakukan pelatihan model menggunakan XGBoost dengan parameter hiper yang telah dioptimalkan. Dari proses tersebut, terbentuk model prediksi penyakit ginjal kronik yang siap digunakan untuk evaluasi tahap berikutnya. Tahap pemrosesan ini sangat menentukan dalam menghasilkan model yang akurat dan dapat dipercaya untuk memprediksi penyakit ginjal kronik, sementara penalaan parameter hiper memastikan bahwa model mengalami peningkatan yang berarti.

3. Evaluasi

Tahap berikutnya adalah tahap evaluasi, tempat performa model prediksi penyakit ginjal kronik ditelaah secara menyeluruh. Pada fase evaluasi ini, digunakan dua alat utama, yaitu matriks kebingungan dan laporan klasifikasi. Matriks kebingungan menyajikan gambaran yang komprehensif mengenai kinerja model melalui pengungkapan jumlah prediksi yang benar maupun yang keliru, sekaligus memperlihatkan pemisahan antara hasil positif dan negatif [21]. Adapun laporan klasifikasi menawarkan metrik penilaian yang lebih detail, misalnya akurasi, presisi, recall, serta skor F1, yang memudahkan penentuan seberapa baik model dapat mengenali penyakit ginjal kronik.



Gambar 3 Confusion Matrix

Matriks kebingungan ini memperlihatkan bagaimana model bekerja saat memisahkan kelas negatif (0) dan kelas positif (1). Angka 53 menandakan banyaknya data negatif yang berhasil diprediksi dengan tepat sebagai negatif (True Negative), sedangkan angka 28 menunjukkan jumlah data positif yang juga berhasil dikenali dengan benar sebagai positif (True Positive). Di sisi lain, hanya ada 1 kasus di mana data yang seharusnya negatif justru diprediksi sebagai positif (False Positive), dan tidak ditemukan data positif yang keliru diprediksi sebagai negatif (False Negative). Temuan tersebut mengindikasikan bahwa model memiliki performa yang sangat baik dalam mengenali kedua kelas, terutama karena tidak adanya kesalahan pada kelas positif.

Berdasarkan confusion matrix tersebut, dapat dilakukan perhitungan metrik evaluasi yang kemudian dirangkum dalam classification report. Total data pengujian adalah 82 data, yang berasal dari penjumlahan seluruh elemen confusion matrix ($53 + 1 + 0 + 28 = 82$). Nilai akurasi dihitung dengan menggunakan rumus berikut:

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

$$accuracy = \frac{28 + 53}{82}$$

$$accuracy = \frac{81}{82}$$

$$accuracy = 0.9878$$

Selanjutnya, untuk kelas 0 (negatif), precision dihitung sebagai:

$$Precision = \frac{TN}{TN + FN}$$

$$Precision = \frac{53}{53 + 0}$$

$$Precision = \frac{53}{53}$$

$$Precision = 1.0000$$

dan recall:

$$Recall = \frac{TN}{TN + FP}$$

$$Recall = \frac{53}{53 + 1}$$

$$Recall = \frac{53}{54}$$

$$Recall = 0.9815$$

Sedangkan untuk kelas 1 (positif), precision adalah:

$$Precision = \frac{TP}{TP + FP}$$

$$Precision = \frac{28}{28 + 1}$$

$$Precision = \frac{28}{29}$$

$$Precision = 0.9655$$

dan recall:

$$Recall = \frac{TP}{TP + FN}$$

$$Recall = \frac{28}{28 + 0}$$

$$Recall = \frac{28}{28}$$

$$Recall = 1.0000$$

Dari nilai precision dan recall tersebut kemudian dihitung F1-score menggunakan rumus:

$$F1 - Score = 2 \times \frac{\text{Precision} \times \text{recall}}{\text{Precision} + \text{recall}}$$

sehingga diperoleh F1-score untuk kelas 0 sebesar 0.9907 dan kelas 1 sebesar 0.9825.

Rata-rata makro (macro average) didapat dengan menyamaratakan nilai dari setiap kelas, sedangkan rata-rata tertimbang memperhitungkan besarnya jumlah data pada masing-masing kelas. Secara keseluruhan, capaian memperlihatkan precision rata-rata makro sebesar 0.9828, recall sebesar 0.9907, dan F1-score sebesar 0.9866. Di sisi lain, precision rata-rata tertimbang tercatat sebesar 0.9882, recall sebesar 0.9878, serta F1-score sebesar 0.9879.

Dari classification report tersebut, dapat ditarik kesimpulan bahwa XGBoost Classifier menunjukkan kinerja yang sangat memuaskan. Keunggulan ini tampak pada perolehan precision, recall, dan F1-score yang tinggi pada kedua kelas, sekaligus menandakan bahwa model menjaga performanya secara seimbang saat mengenali kelas positif maupun kelas negatif. Selain itu, ketiadaan false negative memperlihatkan bahwa model mampu menangkap seluruh kasus positif secara akurat, sehingga berpotensi dimanfaatkan sebagai alat bantu untuk deteksi dini penyakit ginjal kronik.

Temuan penelitian ini selaras dengan studi Nimithap dan rekan-rekan (2026) yang menggunakan algoritma Random Forest untuk memprediksi perkembangan Chronic Kidney Disease berdasarkan data rumah sakit. Studi tersebut melaporkan akurasi sebesar 96,25% dengan ROC-AUC bernilai 1,00, sekaligus memperlihatkan bahwa modelnya mampu melakukan klasifikasi dengan sangat baik dalam membedakan pasien CKD dan non-CKD. Penelitian itu juga menegaskan bahwa hemoglobin, serum kreatinin, packed cell volume, serta specific gravity merupakan atribut yang paling dominan, berdasarkan analisis SHAP [10]. Namun, bila dibandingkan dengan penelitian tersebut, model XGBoost Classifier yang dirancang pada penelitian ini menghasilkan kinerja yang lebih tinggi, yakni akurasi 98,78%, dengan hanya 1 false positive dan 0 false negative. Hasil ini menunjukkan bahwa XGBoost memiliki kemampuan yang sangat baik dalam mengklasifikasikan data penyakit ginjal kronik pada kumpulan data yang digunakan. Meski demikian, kedua penelitian sama-sama menegaskan bahwa penerapan algoritma machine learning dapat menghasilkan performa prediksi yang tinggi dan berpeluang untuk diimplementasikan sebagai sistem pendukung keputusan guna membantu deteksi dini penyakit ginjal kronik.

IV. SIMPULAN

Berdasarkan penelitian yang telah dilakukan, algoritma XGBoost Classifier memperlihatkan kinerja yang sangat kuat saat digunakan untuk melakukan prediksi penyakit ginjal kronik. Dari proses evaluasi, model memperoleh nilai accuracy sebesar 98.78%, dengan confusion matrix yang mencatat 53 True Negative, 28 True Positive, 1 False Positive, dan 0 False Negative. Selain itu, nilai precision, recall, serta F1-score pada kedua kelas juga tergolong tinggi dan tetap seimbang, menandakan bahwa model mampu mengklasifikasikan data dengan tingkat kesalahan yang sangat rendah. Tidak adanya False Negative berarti seluruh pasien yang terindikasi penyakit ginjal kronik berhasil terdeteksi oleh model. Karena itu, algoritma XGBoost Classifier memiliki prospek yang baik untuk dijadikan sistem pendukung keputusan guna membantu deteksi dini penyakit ginjal kronik.

Ucapan Terima Kasih

Puji syukur ke hadirat Allah SWT atas rahmat dan karunia-Nya sehingga penelitian dan penyusunan jurnal ini dapat diselesaikan dengan baik. Penulis mengucapkan terima kasih kepada orang tua dan keluarga atas doa, dukungan, dan motivasi yang diberikan selama proses perkuliahan hingga penyelesaian penelitian ini. Penulis menyadari bahwa penelitian ini masih memiliki keterbatasan. Oleh karena itu, kritik dan saran yang membangun sangat diharapkan untuk penyempurnaan penelitian selanjutnya. Semoga penelitian ini dapat memberikan manfaat bagi pengembangan ilmu pengetahuan dan menjadi referensi bagi penelitian di bidang terkait.

REFERENSI

- [1] W. Yunus, "Algoritma K-Nearest Neighbor Berbasis Particle Swarm Optimization Untuk Prediksi Penyakit Ginjal Kronik," *J. Tek. Elektro CosPhi*, vol. 2, no. 2, pp. 51–55, 2018.
- [2] H. Amalia, "Perbandingan Metode Data Mining Svm Dan Nn Untuk Klasifikasi Penyakit Ginjal Kronis," *Jurnal PILAR Nusa Mandiri*, vol. 14, no. 1, pp. 1–6, 2018.
- [3] F. Sanmarchi, C. Fanconi, D. Golinelli, D. Gori, T. Hernandez-Boussard, and A. Capodici, "Predict, diagnose, and treat chronic kidney disease with machine learning: a systematic literature review," *J. Nephrol.*, vol. 36, no. 4, pp. 1101–1117, 2023.
- [4] T. Arifin and D. Ariesta, "Prediksi Penyakit Ginjal Kronis Menggunakan Algoritma Naive Bayes Classifier Berbasis Particle Swarm Optimization," *Jurnal Tekno Insentif*, vol. 13, no. 1, pp. 26–30, 2019.
- [5] Q. A'yuniyah, E. Tasia, N. Nazira, P. F. Pratama, and M. R. Anugrah, "Implementasi algoritma Naïve Bayes Classifier (NBC) untuk klasifikasi penyakit ginjal kronik," *Jurnal Sistem Komputer dan Informatika (JSON) Hal*, vol. 72, p. 76, 2022.
- [6] MENTERI KESEHATAN REPUBLIK INDONESIA, *KEPUTUSAN MENTERI KESEHATAN REPUBLIK INDONESIA NOMOR HK.01.07/MENKES/1634/2023 TENTANG PEDOMAN NASIONAL PELAYANAN KEDOKTERAN TATA LAKSANA PENYAKIT GINJAL KRONIK*. Indonesia, 2023.
- [7] H. Setiawan, C. Fatichah, and A. Saikhu, "Multilabel classification of student feedback data using BERT and machine learning methods," in *2023 14th International Conference on Information & Communication Technology and System (ICTS)*, IEEE, 2023, pp. 147–152.
- [8] A. Eviyanti, A. Saikhu, and C. Fatichah, "Epileptic Seizure Detection Using Machine Learning and Deep Learning Method," in *2022 IEEE International Conference on Cybernetics and Computational Intelligence (CyberneticsCom)*, IEEE, 2022, pp. 463–468.
- [9] I. R. I. Astutik, U. Indahyanti, and E. Rinata, "Optimization of Stunting Prevention and Reduction through Early Detection Application, Sunting, based on Forward Chaining Inference Machine.: Optimalisasi Pencegahan dan Penurunan Stunting Melalui Aplikasi Deteksi Dini Sunting Berbasis Mesin Inferensi Forward Chaining," *Academia Open*, vol. 8, no. 2, pp. 10–21070, 2023.
- [10] S. Nimithap, S. Bag, S. Elavarasan, P. Ghosh, M. Sathiyamoorthy, and S. T. Gopukumar, "Machine Learning Algorithms for Predicting CKD Progression: A Real-World Hospital Dataset Analysis," *KIDNEYS*, vol. 15, no. 1, pp. 43–54, 2026.
- [11] M. J. Raihan, M. A.-M. Khan, S.-H. Kee, and A.-A. Nahid, "Detection of the chronic kidney disease using XGBoost classifier and explaining the influence of the attributes on the model using SHAP," *Sci. Rep.*, vol. 13, no. 1, p. 6263, 2023.
- [12] I. W. 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 (JTIK)*, vol. 7, no. 2, pp. 417–426, 2020.
- [13] H.-H. Chang, J.-H. Chiang, C.-C. Tsai, and P.-F. Chiu, "Predicting hyperkalemia in patients with advanced chronic kidney disease using the XGBoost model," *BMC Nephrol.*, vol. 24, no. 1, p. 169, 2023.
- [14] V. Wulandari, W. J. Sari, Z. Alfian, L. Legito, and T. Arifianto, "Implementasi Algoritma Naïve Bayes Classifier dan K-Nearest Neighbor untuk Klasifikasi Penyakit Ginjal Kronik: Implementation of Naïve Bayes Classifier and K-Nearest Neighbor Algorithms for Chronic Kidney Disease Classification," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 2, pp. 710–718, 2024.
- [15] W. D McGinnis, C. Siu, A. S, and H. Huang, "Category Encoders: a scikit-learn-contrib package of transformers for encoding categorical data," *The Journal of Open Source Software*, vol. 3, no. 21, 2018, doi: 10.21105/joss.00501.
- [16] Q. A'yuniyah, E. Tasia, N. Nazira, P. F. Pratama, and M. R. Anugrah, "Implementasi algoritma Naïve Bayes Classifier (NBC) untuk klasifikasi penyakit ginjal kronik," *Jurnal Sistem Komputer dan Informatika (JSON) Hal*, vol. 72, p. 76, 2022.
- [17] C. Delrue, S. De Bruyne, and M. M. Speckaert, "Application of machine learning in chronic kidney disease: current status and future prospects," *Biomedicine*, vol. 12, no. 3, p. 568, 2024.
- [18] F. Sanmarchi, C. Fanconi, D. Golinelli, D. Gori, T. Hernandez-Boussard, and A. Capodici, "Predict, diagnose, and treat chronic kidney disease with machine learning: a systematic literature review," *J. Nephrol.*, vol. 36, no. 4, pp. 1101–1117, 2023.
- [19] M. K. Nasution, Rd. R. Saedudin, and V. P. Widartha, "Perbandingan Akurasi Algoritma Naïve Bayes Dan Algoritma Xgboost Pada Klasifikasi Penyakit Diabetes," *e-Proceeding of Engineering*, vol. 8, no. 5, 2021.
- [20] D. A. Anggoro and N. A. Afdallah, "Grid Search CV Implementation in Random Forest Algorithm to Improve Accuracy of Breast Cancer Data," *Int. J. Adv. Sci. Eng. Inf. Technol.*, vol. 12, no. 2, 2022, doi: 10.18517/ijaseit.12.2.15487.

- [21] B. P. Pratiwi, A. S. Handayani, and S. Sarjana, "PENGUKURAN KINERJA SISTEM KUALITAS UDARA DENGAN TEKNOLOGI WSN MENGGUNAKAN CONFUSION MATRIX," *Jurnal Informatika Upgris*, vol. 6, no. 2, 2021, doi: 10.26877/jiu.v6i2.6552.

Conflict of Interest Statement:

The author declares that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.