

**Analisis Prediksi Lama Rawat Inap Pasien Diabetes Tipe 2
Menggunakan Algoritma Decision Tree, Naïve Bayes
Dan K-Nearest Neighbor**

Amelia Tri Ambarwati¹⁾, Wahyu Wijaya Widiyanto²⁾, Frestiany Regina Putri³⁾

^{1,2,3} Politeknik Indonusa Surakarta

^{1,2,3} JL. KH.. Samanhudi No.31, Bumi Laweyan, Surakarta

¹f22033@poltekindonusa.ac.id ²wahyuwijaya@poltekindonusa.ac.id,

³Frestiyay.Putri@poltekindonusa.ac.id

Abstrak

Penelitian ini bertujuan menganalisis kemampuan algoritma *Decision tree*, *Naive bayes*, dan *K-nearest neighbor* (KNN) dalam memprediksi lama rawat inap pasien Diabetes Mellitus Tipe 2. Penelitian dilatarbelakangi oleh meningkatnya jumlah kasus rawat inap dengan diagnosis Diabetes Mellitus Tipe 2 pada periode 2022–2024, yang berpotensi menambah beban pelayanan rumah sakit. Penelitian menerapkan metode deskriptif kuantitatif dengan pendekatan retrospektif yang dilakukan di RSUP dr. Soeradji Tirtonegoro Klaten. Sumber data yang digunakan berupa rekam medis elektronik pasien rawat inap yang memiliki diagnosis Diabetes Mellitus Tipe 2 selama tahun 2022–2024. Analisis data dilaksanakan berdasarkan kerangka kerja *Knowledge Discovery in Database* (KDD), yang terdiri atas *tahapan data selection, preprocessing, transformation, data mining*, dan *interpretation/evaluation*. Dalam tahap data mining, model klasifikasi dibangun menggunakan algoritma *Decision tree*, *Naive bayes*, dan KNN. Temuan penelitian mengindikasikan bahwa ketiga algoritma tersebut memiliki kemampuan untuk melakukan klasifikasi dalam prediksi lama rawat inap pasien Diabetes Mellitus Tipe 2. Hasil pengujian menunjukkan bahwa algoritma *Naive bayes* memberikan performa terbaik dengan nilai *accuracy* sebesar 0,525. Kinerja tersebut melampaui *Decision tree* yang memperoleh nilai *accuracy* 0,501 serta KNN dengan nilai *accuracy* 0,438. Berdasarkan hasil tersebut, *Naive bayes* ditetapkan sebagai algoritma yang paling unggul dalam memprediksi lama rawat inap pasien pada penelitian ini.

Kata kunci: Diabetes Melitus Tipe 2, Prediksi Lama Rawat Inap, *Decision tree*, *Naïve Bayes*, *K-nearest neighbor*

Abstract

This study aims to analyze the ability of Decision tree, Naive bayes, and K-nearest neighbor (KNN) algorithms in predicting the length of stay of patients with Type 2 Diabetes Mellitus. The research is motivated by the increasing number of hospitalizations with a diagnosis of Type 2 Diabetes Mellitus in the period 2022–2024, which has the potential to increase the burden on hospital services. The study applies a quantitative descriptive method with a retrospective approach conducted at Dr. Soeradji Tirtonegoro General Hospital, Klaten. The data source used is the electronic medical records of inpatients diagnosed with Type 2 Diabetes Mellitus during 2022–2024. Data analysis is carried out based on the Knowledge Discovery in Database (KDD) framework, which consists of data selection, preprocessing, transformation, data mining, and interpretation/evaluation stages. In the data mining stage, a classification model is built using Decision tree, Naive bayes, and KNN algorithms. The research findings indicate that the three algorithms have the ability to classify in predicting the length of stay of patients with Type 2 Diabetes Mellitus. The test results show that the Naive bayes algorithm provides the best performance with an accuracy value of 0.525. This performance surpasses Decision tree which obtained an accuracy value of 0.501 and KNN with an accuracy value of 0.438. Based on these results, Naive bayes is determined to be the most superior algorithm in predicting the length of stay of patients in this study.

Keywords: Type 2 Diabetes Mellitus, Length of Stay Prediction, Decision tree, Naive bayes, K-nearest neighbor

1. PENDAHULUAN

Pesatnya perkembangan teknologi informasi memberikan pengaruh yang luas terhadap berbagai sektor, salah satunya adalah sektor kesehatan, yang turut mengalami transformasi dalam berbagai aspek pelayanan dan pengelolaannya. Salah satu bentuk transformasi yang terjadi di bidang kesehatan adalah penerapan digitalisasi data medis melalui penggunaan Rekam Medis Elektronik (RME). Implementasi Rekam Medis Elektronik (RME) telah memiliki landasan hukum yang ditetapkan melalui Peraturan Menteri Kesehatan Republik Indonesia Nomor 24 Tahun 2022 tentang Rekam Medis Elektronik (Kementrian Kesehatan RI, 2022) penerapan RME memungkinkan data pasien terdokumentasi secara sistematis, terstruktur, dan mudah diakses, sehingga mendukung efektivitas pelayanan medis serta pengambilan keputusan klinis yang lebih tepat.

Seiring meningkatnya volume data medis yang tersimpan dalam RME, muncul tantangan baru terkait pemanfaatan data secara optimal. Data yang melimpah belum tentu memberikan manfaat maksimal apabila tidak diolah dengan metode yang tepat. Di sinilah peran teknologi analitik seperti data mining menjadi sangat penting (Romlah & Laiman, 2022); (Limbong & Afriani, 2023). Data mining merupakan proses analitis untuk menemukan pola dan pengetahuan yang berguna dari kumpulan data besar, sehingga dapat mendukung pengambilan keputusan secara prediktif dan proaktif (Nugroho et al., 2022). Dalam pelayanan rumah sakit, salah satu indikator penting yang perlu dikelola dengan baik adalah lama rawat inap pasien. Lama rawat inap yang terlalu panjang dapat menyebabkan penumpukan pasien, pemborosan sumber daya, serta peningkatan biaya operasional (Giusman & Nurwahyuni, 2022).

Sebaliknya, masa rawat inap yang terlalu singkat tanpa kesiapan pasien juga berisiko terhadap keselamatan dan kualitas perawatan. Seperti penelitian dari (Syahputra, 2022), kemampuan memprediksi lama rawat inap sejak awal pasien masuk rumah sakit menjadi menjadi hal yang perlu dipertimbangkan bagi manajemen rumah sakit sehingga nantinya memiliki strategi untuk mempertimbangkan kebutuhan bed hingga SDM dalam usaha meningkatkan efektivitas pelayanan.

Berdasarkan data rawat inap di RSUP dr. Soeradji Tirtonegoro Klaten selama periode 2022– 2024, terdapat beberapa diagnosa penyakit yang cukup dominan dalam menyebabkan pasien menjalani perawatan di rumah sakit. Data menunjukkan bahwa kasus Diabetes Mellitus (DM) Tipe 2 yang diklasifikasi dalam kode E11.9 mengalami peningkatan yang signifikan dibandingkan dengan diagnosa lainnya, seperti dispepsia dan hipertensi. Jumlah pasien rawat inap dengan diagnosa Diabetes Mellitus (DM) Tipe 2 tercatat sebanyak 237 kasus pada tahun 2022, meningkat menjadi 437 kasus pada tahun 2023, dan kembali meningkat menjadi 461 kasus pada tahun 2024.

Peningkatan tersebut menunjukkan bahwa Diabetes Mellitus (DM) Tipe 2 merupakan salah satu permasalahan kesehatan yang cukup serius dan berpotensi meningkatkan beban pelayanan rawat inap di rumah sakit. Pasien dengan diagnosa Diabetes Mellitus (DM) Tipe 2 memiliki variasi lama rawat inap yang berbeda-beda, yang dipengaruhi oleh kondisi klinis, faktor komorbiditas, serta respons terhadap terapi yang diberikan. Variasi durasi perawatan tersebut menuntut adanya suatu pendekatan analisis yang mampu memprediksi lama rawat inap secara lebih sistematis dan terukur.

Dalam penelitian ini, lama rawat inap diklasifikasikan ke dalam tiga kategori, yaitu sebentar (9 hari) (Luxiarti et al., 2022). Salah satu metode yang dapat dimanfaatkan untuk memperkirakan lama rawat inap pasien adalah teknik klasifikasi dalam data mining. Berbagai penelitian terdahulu menunjukkan bahwa algoritma klasifikasi seperti *Decision tree*, *Naïve Bayes*, dan *K-nearest neighbor* (KNN) memiliki kinerja yang baik serta banyak diterapkan dalam proses prediksi dan pengambilan keputusan berbasis data. Masing-masing algoritma memiliki karakteristik yang berbeda dalam membangun model prediksi, seperti struktur pohon keputusan pada *Decision tree*, pendekatan probabilistik pada *Naïve Bayes*, serta perhitungan kedekatan jarak pada KNN.

Perbedaan karakteristik tersebut menyebabkan perlunya analisis komparatif untuk menentukan algoritma dengan tingkat akurasi terbaik dalam memprediksi lama rawat inap pasien (Biyantoro & Prasetyo, 2024). Berdasarkan latar belakang tersebut, dibutuhkan metode analisis yang dapat memanfaatkan data pasien rawat inap secara

optimal dan sistematis guna memprediksi lama perawatan pasien. Oleh karena itu, penelitian ini dilakukan dengan tujuan untuk mengevaluasi dan membandingkan tingkat ketepatan algoritma *Decision tree*, *Naive bayes*, dan *K-nearest neighbor* (KNN) dalam melakukan prediksi lama rawat inap pasien Diabetes Mellitus Tipe 2 (E11.9). Data yang digunakan berasal dari pasien rawat inap di RSUP dr. Soeradji Tirtonegoro Klaten selama periode 2022–2024. Penelitian ini diharapkan dapat menghasilkan model prediksi dengan tingkat akurasi terbaik sehingga dapat mendukung proses perencanaan dan pengambilan keputusan dalam manajemen pelayanan rumah sakit.

2. TINJAUAN PUSTAKA

a. Rekam Medis Elektronik

Menurut Peraturan Menteri Kesehatan (Permenkes) Nomor 24 Tahun 2022, Rekam Medis Elektronik (RME) adalah bentuk rekam medis yang disusun, dikelola, dan disimpan dengan menggunakan sistem elektronik. RME memungkinkan pencatatan informasi kesehatan secara lebih cepat, akurat, dan terintegrasi, sehingga dapat meningkatkan efisiensi pelayanan serta mempermudah akses data bagi tenaga kesehatan yang berwenang. Dokumentasi elektronik ini dirancang untuk mendukung penyelenggaraan rekam medis secara modern sesuai dengan perkembangan teknologi informasi di bidang kesehatan (Kemenkes RI, 2022).

b. Lama Rawat Inap

Lama hari rawat merupakan durasi yang menunjukkan berapa lama seorang pasien menjalani perawatan di rumah sakit, terhitung sejak hari pertama masuk hingga pasien dinyatakan boleh pulang. Indikator ini sering dijadikan tolok ukur oleh rumah sakit dalam menilai kualitas pelayanan yang diberikan, karena lama hari rawat yang optimal mencerminkan efektivitas penanganan medis, ketepatan diagnosis, dan efisiensi tindakan terapeutik yang diberikan kepada pasien. Selain berfungsi sebagai alat evaluasi mutu layanan, lama hari rawat juga menjadi dasar bagi rumah sakit dalam merencanakan kebutuhan sumber

daya, seperti tenaga medis, fasilitas ruang perawatan, serta obat-obatan dan peralatan penunjang lainnya. Dengan memantau dan menganalisis lama hari rawat secara berkala, rumah sakit dapat mengidentifikasi potensi kendala dalam proses perawatan, melakukan perbaikan strategi pelayanan, serta meningkatkan kepuasan pasien melalui penyediaan layanan kesehatan yang lebih cepat, tepat, dan berkualitas Ramadhani et al. (2025).

c. Diabetes Mellitus Tipe 2

Non-insulin-dependent diabetes mellitus without complications adalah diabetes melitus tipe 2 tanpa komplikasi, yaitu kondisi kronis yang ditandai dengan meningkatnya kadar gula darah akibat tubuh tidak mampu menggunakan insulin secara efektif (resistensi insulin) atau produksi insulin yang tidak mencukupi, namun belum menimbulkan komplikasi pada organ tubuh seperti ginjal, mata, saraf, atau jantung. Pada kondisi ini, pasien umumnya tidak bergantung pada terapi insulin sebagai pengobatan utama dan 10 dapat dikontrol melalui pola makan, olahraga, dan obat antidiabetes oral. Dalam sistem klasifikasi penyakit ICD-10, kondisi ini dikodekan sebagai E11.9 dan termasuk dalam kategori diabetes melitus tipe 2 tanpa komplikasi

d. Prediksi

Prediksi (peramalan) merupakan suatu upaya untuk memperkirakan atau meramalkan peristiwa yang akan terjadi pada masa mendatang dengan memanfaatkan informasi dari peristiwa pada masa lampau. Proses prediksi dilakukan melalui pengumpulan data historis (*history data*), yaitu data yang telah terjadi sebelumnya dan dianggap relevan dengan kondisi atau variabel yang ingin diramalkan. Data historis tersebut kemudian dianalisis menggunakan metode statistik atau pendekatan tertentu untuk mengidentifikasi pola, kecenderungan, atau hubungan antar variabel. Berdasarkan pola atau corak yang ditemukan, dilakukan perhitungan lebih lanjut untuk menghasilkan nilai prediksi yang dapat digunakan sebagai acuan atau gambaran mengenai keadaan di masa depan. Dengan demikian, prediksi

berperan penting dalam proses perencanaan, pengambilan keputusan, serta penentuan strategi yang lebih tepat dan akurat (Rahayu, 2023)

e. Data Mining

Data mining adalah suatu proses penggalian informasi penting dari suatu data besar. Data mining juga dapat didefinisikan sebagai proses penemuan pola, hubungan, atau informasi berguna dari sejumlah besar data yang sering kali tersimpan dalam basis data (Yunardi & Dina, 2022)

f. Decision tree

Decision tree adalah salah satu metode klasifikasi yang digunakan untuk memprediksi suatu pola dari dataset, dengan memanfaatkan hubungan atau relasi antarvariabel untuk menemukan target. Metode ini berbentuk diagram alir top down (dari atas ke bawah) yang menyerupai akar, memiliki batang akar, percabangan dan daun. Masing masing variabel memiliki dua kemungkinan (*probability*) nilai yang bersifat kategorial seperti “Yes” atau “No”. sehingga, setiap atribut merupakan nilai pasti dan tidak memiliki nilai yang berada di antaranya (Yunardi & Dina, 2022)

g. Naive bayes

Naive Bayes merupakan salah satu metode klasifikasi probabilistik yang didasarkan pada penerapan Teorema Bayes. Metode ini bekerja dengan menggunakan asumsi independensi yang kuat antar setiap fitur atau variabel yang digunakan dalam proses klasifikasi. Artinya, setiap fitur dianggap tidak saling bergantung satu sama lain dalam menentukan probabilitas suatu kelas, meskipun pada kenyataannya hubungan antar fitur mungkin saja ada. Salah satu keunggulan utama dari metode *Naive Bayes* adalah kesederhanaannya dalam proses perhitungan serta efisiensinya dalam penggunaan data. Dengan karakteristik tersebut, *Naive Bayes* sering digunakan dalam berbagai permasalahan klasifikasi karena mampu memberikan hasil yang cukup baik meskipun dengan sumber data yang terbatas. Selain itu, metode ini juga memiliki kemampuan

komputasi yang cepat sehingga efektif diterapkan pada dataset berukuran besar (Bilqisth et al., 2022).

h. K-nearest neighbor

Metode *K-nearest neighbor* (KNN) merupakan salah satu algoritma dalam *machine learning* yang digunakan untuk melakukan klasifikasi berdasarkan tingkat kedekatan atau jarak antara suatu data dengan data lainnya. Algoritma ini termasuk dalam kategori *supervised learning*, yaitu metode pembelajaran yang menggunakan data latih yang telah memiliki label atau kategori.

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan deskriptif kuantitatif dengan studi retrospektif. Penelitian dilakukan di RSUP dr. Soeradji Tirtonegoro Klaten dengan objek penelitiannya adalah rekam medis elektronik pasien poli penyakit dalam, periode 2022 – 2024. Pada tahap awal data dikumpulkan secara retrospektif, lalu dilakukan analisis data menggunakan metode KDD yang diawali dengan *data selection*, *preprocessing*, *Transformation*, *data mining* menggunakan algoritma *Decision tree*, *Naive Bayes* dan KNN sampai dilakukan *interpretation/evaluation* (Azhar et al., 2022).

Objek pada penelitian ini menggunakan rekam medis elektronik (RME) pasien rawat inap yang memiliki diagnosa utama Diabetes Melitus Tipe 2 atau *Non-insulin-dependent diabetes mellitus without complications* (E11.9) di RSUP dr. Soeradji Tirtonegoro Klaten. Data tersebut mencakup seluruh dokumen RME yang tercatat sepanjang tahun 2022 – 2024 dengan jumlah total 1.135 dokumen RME. Pengumpulan data dilakukan dengan pendekatan studi retrospektif, data diambil dari sistem informasi rumah sakit yang diekspor dalam bentuk *excel*. Variabel yang digunakan adalah usia, jenis kelamin, diagnosa sekunder, tindakan, cara bayar, cara masuk dan LOS (*Lenght Of Stay*). Metode *Knowledge Discovery in Database* (KDD) digunakan sebagai kerangka kerja dalam analisis data penelitian ini. Proses tersebut mencakup tahapan *data selection*, *preprocessing*, *transformation*, *data mining*, hingga *interpretation/evaluation*. Pada proses data mining, klasifikasi data dilakukan dengan memanfaatkan tiga algoritma, yaitu *Decision tree*, *Naive bayes*, serta *K-Nearest Neighbor*

(KNN), untuk menghasilkan model prediksi yang akan dievaluasi kinerjanya. Penelitian ini juga menerapkan pendekatan analisis komparatif untuk membandingkan performa ketiga metode klasifikasi tersebut. Proses analisis dilakukan dengan menerapkan masing-masing algoritma pada dataset yang sama, kemudian hasil klasifikasi dievaluasi menggunakan metrik akurasi. Perbandingan nilai akurasi dari setiap metode digunakan untuk menentukan algoritma dengan performa terbaik dalam melakukan klasifikasi pada data penelitian

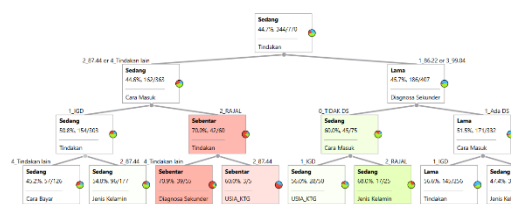
4. HASIL DAN EMBahasan

a. Hasil

Proses analisis data diawali dengan tahap *data selection*, yaitu pemilihan variabel yang digunakan dalam penelitian meliputi usia, jenis kelamin, diagnosis sekunder, tindakan, cara bayar, cara masuk pasien, dan lama rawat inap (*Length of Stay/LOS*). Setelah itu, dataset melalui tahap *preprocessing* menggunakan *Power BI* karena sebelum dilakukan analisis dan pemodelan, kualitas data perlu ditingkatkan melalui tahap *preprocessing*.

Tahapan ini mencakup pembersihan data (*data cleaning*), Tahap *preprocessing* meliputi pembersihan data (*data cleaning*) berupa pemeriksaan *missing value*, penghapusan data duplikat, penyesuaian format penulisan, serta pengecekan konsistensi antar atribut. Selanjutnya dilakukan transformasi data, yaitu mengubah data kategori menjadi numerik (*encoding*), menyeragamkan format penulisan data, dan melakukan pengelompokan data tertentu agar dataset lebih terstruktur dan siap digunakan dalam proses pemodelan penelitian.

1) Hasil Prediksi *Decision tree*



Gambar 1. Visualisasi *Decision Tree*

Keterangan:

Warna biru: Kategori LOS Lama (≥ 9 hari)

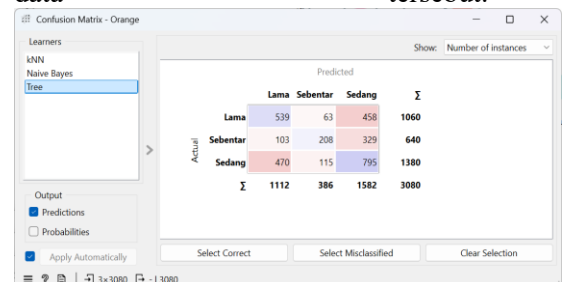
Warna hijau: Kategori LOS Sedang (≥ 5 hari)

Warna Merah: Kategori LOS Sebentar (<5 hari)

Semakin pekat warna node menunjukkan tingkat dominasi salah satu kategori.

Gambar 1 menunjukkan visualisasi hasil pemodelan klasifikasi dengan algoritma *Decision tree* yang dihasilkan melalui aplikasi *Orange Data Mining*. Berdasarkan hasil pembentukan pohon keputusan, variabel tindakan menjadi node akar (*root node*) yang berfungsi sebagai atribut utama dalam proses klasifikasi. Selanjutnya, proses pemisahan data dilakukan menggunakan beberapa variabel lain seperti jenis kelamin, usia, cara masuk, diagnosis, dan lama perawatan atau LOS.

Pada pohon keputusan, masing-masing cabang menunjukkan proses pengambilan keputusan yang didasarkan pada nilai atribut tertentu sebagai dasar klasifikasi. sedangkan node akhir (*leaf node*) menunjukkan hasil klasifikasi yang ditunjukkan model. Warna pada *node* merepresentasikan kategori kelas dominan yang diprediksi oleh model. Di mana semakin pekat warna *node* menunjukkan tingkat dominasi kelas yang semakin tinggi pada kelompok data tersebut.



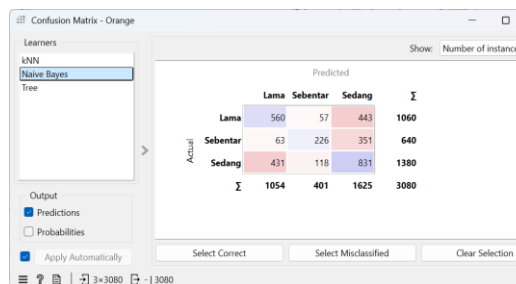
Gambar 2. *Confusion matrix Decision Tree*

Gambar 2 menampilkan hasil analisis *confusion matrix* algoritma *decision tree* pada aplikasi *Orange Data Mining*. Berdasarkan hasil *confusion matrix*, total data *testing* yang digunakan sebanyak 3080 data pasien, terdiri dari 1060 data kategori “Lama”, 640 data kategori “Sebentar”, dan 1380 data

kategori “Sedang”. *Confusion matrix* digunakan untuk mengukur tingkat kesesuaian antara kelas aktual dengan hasil prediksi yang dihasilkan oleh model *decision tree*. Pada kategori Lama, model mampu memprediksi dengan benar sebanyak 539 data. Namun, masih terdapat kesalahan prediksi, yaitu 63 data diprediksi sebagai kategori Sebentar dan 458 data diprediksi sebagai kategori Sedang. Pada kategori Sebentar, model berhasil memprediksi dengan benar sebanyak 208 data, sedangkan 103 data salah diprediksi menjadi kategori Lama dan 329 data diprediksi sebagai kategori Sedang. Sementara itu, pada kategori Sedang, model berhasil memberikan prediksi benar sebanyak 795 data. Akan tetapi, terdapat 470 data yang salah diprediksi sebagai kategori Lama dan 115 data diprediksi sebagai kategori Sebentar. Secara keseluruhan, model berhasil memprediksi dengan tepat sebanyak 1542 data dari total 3080 data pengujian.

2) Hasil Prediksi *Naive bayes*

Gambar 3 di bawah menampilkan hasil analisis *confusion matrix* algoritma *naive bayes* pada aplikasi Orange Data Mining. Berdasarkan hasil *confusion matrix*, total data testing yang digunakan sebanyak 3080 data pasien, yang terdiri dari 1060 data kategori “Lama”, 640 data kategori “Sebentar”, dan 1380 data kategori “Sedang”. *Confusion matrix* digunakan untuk mengetahui tingkat kesesuaian antara kelas aktual dengan hasil prediksi yang dihasilkan oleh model *naive bayes*.



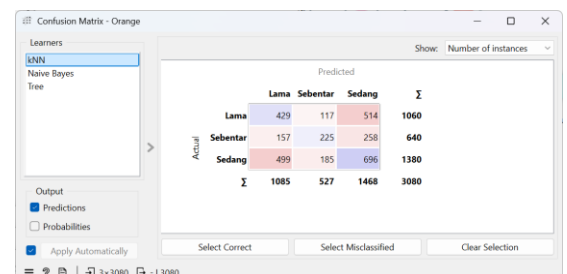
Gambar 3. *Confusion Matrix Naive Bayes*

Pada kategori Lama, terdapat 1060 data aktual, di mana sebanyak 560 data berhasil diprediksi dengan benar

sebagai kategori “Lama”. Namun, masih terdapat kesalahan prediksi, yaitu 57 data diprediksi sebagai kategori “Sebentar” dan 443 data diprediksi sebagai kategori “Sedang”. Pada kategori Sebentar, dari total 640 data aktual, model mampu memprediksi dengan benar sebanyak 226 data sebagai kategori “Sebentar”, sedangkan 63 data salah diprediksi menjadi kategori “Lama” dan 351 data diprediksi menjadi kategori “Sedang”. Sementara itu, pada kategori Sedang terdapat 1380 data aktual. Dari jumlah tersebut, sebanyak 831 data berhasil diprediksi dengan benar sebagai kategori “Sedang”. Akan tetapi, model masih melakukan kesalahan prediksi terhadap 431 data yang diprediksi menjadi kategori “Lama” dan 118 data diprediksi sebagai kategori “Sebentar”. Secara keseluruhan, model berhasil memprediksi dengan benar sebanyak 1617 data dari total 3080 data pengujian.

3) Hasil Prediksi *K-nearest neighbor*

Gambar 4 di bawah menampilkan hasil analisis *confusion matrix* algoritma *K-nearest neighbor* (KNN) pada aplikasi Orange Data Mining. Berdasarkan hasil *confusion matrix*, total data testing yang digunakan sebanyak 3080 data pasien, terdiri dari 1060 data kategori “Lama”, 640 data kategori “Sebentar”, dan 1380 data kategori “Sedang”. *Confusion matrix* digunakan untuk mengetahui tingkat kesesuaian antara kelas aktual dengan hasil prediksi yang dihasilkan oleh model KNN.



Gambar 4. *Confusion Matrix K-Nearest Neighbor*

Pada kategori Lama, model berhasil memprediksi dengan benar sebanyak 429 data sebagai kategori “Lama”. Namun, masih terdapat kesalahan prediksi, yaitu 117 data diprediksi

sebagai kategori “Sebentar” dan 514 data diprediksi sebagai kategori “Sedang”. Pada kategori Sebentar, dari total 640 data aktual, hanya 225 data yang berhasil diprediksi dengan benar, sedangkan 157 data salah diprediksi menjadi kategori “Lama” dan 258 data diprediksi sebagai kategori “Sedang”. Sementara itu, pada kategori Sedang, model mampu memprediksi dengan benar sebanyak 696 data dari total 1380 data aktual. Akan tetapi, masih ditemukan kesalahan prediksi, yaitu 499 data diprediksi menjadi kategori “Lama” dan 185 data diprediksi sebagai kategori “Sebentar”. Secara keseluruhan, model berhasil memberikan prediksi yang benar sebanyak 1350 data dari total 3080 data pengujian.

4) Perbandingan Evaluasi Model

Model	Evaluasi					
	AUC	CA	F1	Proc	Recall	MCC
<i>Decision tree</i>	0,631	0,501	0,496	0,504	0,501	0,196
<i>Naive bayes</i>	0,663	0,525	0,520	0,529	0,525	0,235
<i>K-nearest neighbor</i>	0,598	0,438	0,437	0,437	0,438	0,108

Gambar 5. Perbandingan Evaluasi Model

Gambar 5 menyajikan hasil perbandingan evaluasi dari tiga algoritma klasifikasi yang diterapkan dalam prediksi lama rawat inap pasien, yaitu *Decision tree*, *Naive bayes*, dan *K-nearest neighbor* (KNN). Evaluasi kinerja model dilakukan menggunakan sejumlah metrik, antara lain AUC (*Area Under Curve*), CA (*Classification Accuracy*), F1-Score, *Precision*, *Recall*, serta MCC (*Matthews Correlation Coefficient*). AUC digunakan sebagai indikator untuk menilai kemampuan model dalam membedakan antar kelas data yang berbeda. Sementara itu, CA (*Classification Accuracy*) menunjukkan tingkat akurasi keseluruhan dari hasil prediksi model. F1-Score berperan dalam mengevaluasi keseimbangan antara *Precision* dan *Recall*, khususnya pada kondisi data yang memiliki ketidakseimbangan kelas. *Precision* sendiri digunakan untuk mengukur ketepatan prediksi terhadap suatu kelas tertentu, sedangkan *Recall*

mengukur kemampuan model dalam menangkap seluruh data yang sebenarnya termasuk dalam kelas tersebut. Selain itu, MCC digunakan untuk mengevaluasi kualitas klasifikasi secara menyeluruh dengan mempertimbangkan seluruh hasil prediksi benar maupun salah sehingga menghasilkan penilaian performa yang lebih seimbang. Berdasarkan tabel 4.2, algoritma *Naive bayes* memperoleh hasil evaluasi terbaik dibandingkan algoritma lainnya. Hal ini dapat dilihat dari nilai AUC sebesar 0,663, nilai akurasi (CA) sebesar 0,525, nilai F1 sebesar 0,520, nilai *Precision* sebesar 0,529, nilai *Recall* sebesar 0,525, dan nilai MCC sebesar 0,235. Metode *Decision tree* berada pada urutan kedua dengan nilai AUC sebesar 0,631 dan nilai akurasi (CA) sebesar 0,501. Selain itu, metode ini memperoleh nilai F1 sebesar 0,496, *Precision* sebesar 0,504, *Recall* sebesar 0,501, dan MCC sebesar 0,196. Metode *K-nearest neighbor* (K-NN) memiliki hasil evaluasi paling rendah dibandingkan kedua metode lainnya. Nilai AUC yang diperoleh sebesar 0,598 dengan nilai akurasi (CA) sebesar 0,438, nilai F1 sebesar 0,437, *Precision* sebesar 0,437 dan *Recall* sebesar 0,438 dan nilai MCC sebesar 0,108.

b. Pembahasan

1) Algoritma *Decision tree*

Berdasarkan hasil evaluasi yang disajikan pada Gambar 5, model tersebut menunjukkan kemampuan yang relatif baik dalam mengklasifikasikan kategori lama rawat inap pasien. Nilai akurasi (CA) sebesar 0,501 menunjukkan bahwa model mampu menghasilkan prediksi yang benar pada sekitar setengah dari data uji. Nilai F1-Score sebesar 0,496 menunjukkan keseimbangan yang cukup antara akurasi prediksi dan kemampuan model untuk mengidentifikasi data yang relevan. Nilai *Precision* sebesar 0,504 menunjukkan bahwa akurasi prediksi model berada dalam kategori cukup akurat, sedangkan nilai *Recall* sebesar 0,501 menunjukkan bahwa kemampuan model dalam mengenali data aktual masih berada pada tingkat sedang.

Selain itu, nilai MCC sebesar 0,196 menunjukkan bahwa hubungan antara hasil prediksi dengan data aktual masih relatif lemah, namun tetap lebih baik dibandingkan prediksi secara acak.

Algoritma *decision tree* bekerja dengan membentuk struktur pohon keputusan berdasarkan atribut yang memiliki nilai informasi tertinggi. Metode ini memiliki keunggulan karena mudah dipahami dan mampu menghasilkan aturan klasifikasi secara jelas. Hasil penelitian ini sejalan dengan penelitian yang dilakukan (Sanusi et al., 2023) yang menjelaskan bahwa *decision tree* cukup efektif digunakan pada data kesehatan karena mampu mengenali pola atribut pasien secara sistematis. Akan tetapi, metode ini juga memiliki kelemahan yaitu rentan mengalami *overfitting* apabila data memiliki pola yang kompleks atau distribusi kelas yang tidak seimbang. Hal tersebut terlihat pada penelitian ini dimana nilai evaluasi masih berada pada kategori sedang.

Hasil metode *decision tree* dipengaruhi oleh karakteristik data pasien Diabetes Mellitus Tipe 2 yang memiliki variasi kondisi medis cukup kompleks. Selain itu, hubungan antar variabel kemungkinan tidak bersifat sederhana sehingga sulit dipisahkan hanya berdasarkan aturan pada pohon keputusan. Distribusi kategori lama rawat inap yang tidak merata juga dapat memengaruhi performa model, misalnya jumlah pasien pada kategori Sedang lebih banyak dibandingkan kategori Sebesar atau Lama. Kondisi tersebut dikenal sebagai *imbalanced data* atau data tidak seimbang.

Akibatnya, model *decision tree* cenderung lebih mudah mempelajari pola pada kategori dengan jumlah data lebih banyak, sedangkan kategori dengan jumlah data yang lebih sedikit menjadi lebih sulit dikenali. Hal ini menyebabkan hasil prediksi pada beberapa kategori menjadi kurang optimal atau kurang akurat (Fazriansyah et al., 2025).

2) Algoritma *Naive bayes*

Hasil pengujian yang disajikan pada Gambar 5 menunjukkan bahwa

algoritma *Naive bayes* berkinerja lebih baik dalam memprediksi lama rawat inap pasien dengan Diabetes Melitus tipe 2 dibandingkan metode lain yang digunakan. Nilai AUC sebesar 0,663 menunjukkan kemampuan model untuk membedakan antara kelas lama rawat inap. Lebih lanjut, nilai akurasi sebesar 0,525 menunjukkan bahwa model mampu memprediksi dengan benar lebih dari 50% data uji. Selain itu, skor F1 sebesar 0,520 menunjukkan keseimbangan yang relatif baik antara presisi dan *recall*. Presisi sebesar 0,529 menunjukkan bahwa variasi model cukup akurat dalam mengklasifikasikan data ke dalam kelas-kelas tertentu, sedangkan *recall* sebesar 0,525 menunjukkan kemampuan model untuk secara konsisten mengidentifikasi data aktual. Lebih lanjut, nilai MCC sebesar 0,235 mencerminkan perbedaan positif antara hasil prediksi dan nilai aktual, yang lebih unggul dibandingkan metode lain. Oleh karena itu, *Naive bayes* dianggap sebagai metode dengan stabilitas terbaik dalam proses klasifikasi dalam penelitian ini.

Algoritma *Naive bayes* adalah metode klasifikasi berdasarkan pendekatan probabilistik yang berlandaskan Teorema Bayes, yang mengasumsikan bahwa setiap atribut bersifat independen. Pendekatan ini banyak digunakan di bidang medis karena dianggap efisien dalam memproses dataset besar dan memiliki waktu komputasi yang relatif cepat. Hasil penelitian ini menguatkan temuan sebelumnya oleh (Widodo & Swedia, 2024) pada tahun 2024, yang membandingkan kinerja algoritma *K-nearest neighbor* (KNN) dan *Naive bayes* dalam diagnosis diabetes mellitus. Penelitian tersebut menunjukkan bahwa *Naive bayes* menunjukkan akurasi yang lebih unggul, yaitu 74,7%, dibandingkan dengan KNN, yang hanya mencapai 68,6%. Selain itu, penelitian yang dilakukan oleh (Nasien et al., 2024) pada tahun 2023 juga menyatakan bahwa algoritma *Naive bayes* menunjukkan kinerja optimal dalam klasifikasi diabetes dibandingkan dengan KNN dan Regresi Logistik,

dengan tingkat akurasi 79%.

Hasil dari kedua penelitian terdahulu tersebut mendukung temuan dalam penelitian ini bahwa algoritma *naive bayes* memiliki kemampuan yang lebih baik dalam melakukan klasifikasi data medis, meskipun performanya dalam memprediksi lama rawat inap masih belum optimal. Metode *naive bayes* memperoleh hasil terbaik karena data pasien Diabetes Mellitus Tipe 2 memiliki pola probabilitas yang cukup jelas antar variabel, sehingga hubungan antar atribut pasien dapat dipelajari dengan baik melalui pendekatan probabilistik. Meskipun demikian, hasil evaluasi yang masih berada dalam kategori sedang mengindikasikan bahwa model belum bekerja secara maksimal. Kondisi tersebut kemungkinan dipengaruhi oleh adanya keterkaitan antar atribut serta distribusi data yang tidak seimbang pada setiap kategori lama rawat inap pasien. Namun secara keseluruhan, metode *naive bayes* dinilai paling efektif pada penelitian ini karena menghasilkan performa evaluasi tertinggi dibandingkan metode lainnya.

3) Algoritma *K-nearest neighbor*

Hasil evaluasi pada Gambar 5 mengindikasikan bahwa algoritma *K-nearest neighbor* (KNN) masih belum mampu memberikan hasil prediksi yang baik terhadap lama rawat inap pasien Diabetes Mellitus Tipe 2. Performa yang dihasilkan masih belum optimal jika dibandingkan dengan algoritma *naive bayes* maupun *decision tree*. Nilai AUC sebesar 0,598 menunjukkan bahwa kemampuan model dalam membedakan kategori lama rawat inap masih tergolong rendah. Nilai akurasi sebesar 0,438 juga menunjukkan bahwa model belum mampu menghasilkan prediksi yang baik. Selain itu, nilai *F1-Score* sebesar 0,437 menandakan bahwa keseimbangan antara *Precision* dan *Recall* masih rendah. Nilai *Precision* sebesar 0,437 menunjukkan bahwa tingkat ketepatan prediksi model masih kurang baik, sedangkan nilai *Recall* sebesar 0,438 menunjukkan

kemampuan model dalam mengenali data aktual juga masih rendah. Di samping itu, nilai MCC sebesar 0,108 menunjukkan bahwa hubungan antara hasil prediksi dengan data aktual sangat lemah, sehingga model dinilai kurang stabil dalam melakukan proses klasifikasi.

Algoritma KNN bekerja dengan menentukan kelas suatu data berdasarkan kedekatan jarak terhadap sejumlah tetangga terdekat. Metode ini cukup sederhana dan efektif pada data dengan pola yang jelas. Namun, KNN sangat dipengaruhi oleh jumlah tetangga (*neighbor*), skala data, dan distribusi kelas. Hasil penelitian ini sejalan dengan penelitian yang dilakukan oleh (Hunafa & Hermawan, 2023) menjelaskan bahwa performa algoritma *K-nearest neighbor* (KNN) cenderung tidak selalu menjadi metode terbaik dibandingkan algoritma lain seperti *Naive bayes* maupun *Decision tree*. Hal tersebut disebabkan karena KNN merupakan metode berbasis jarak (*distance-based*) yang sangat bergantung pada kedekatan antar data. Ketika data memiliki kompleksitas tinggi, banyak atribut, maupun distribusi kelas yang tidak seimbang, kinerja KNN dapat menurun karena proses penentuan tetangga terdekat menjadi kurang optimal. Sebaliknya, algoritma seperti *Naive bayes* dan *Decision tree* cenderung lebih mampu menangani pola distribusi data tertentu sehingga menghasilkan performa klasifikasi yang lebih baik pada beberapa kasus diabetes.

Rendahnya performa algoritma KNN dapat dipengaruhi oleh beberapa faktor, salah satunya karakteristik data medis yang kompleks serta adanya kemiripan antar kategori lama rawat inap. Pasien dengan kondisi klinis yang serupa tidak selalu memiliki durasi rawat inap yang sama karena dipengaruhi oleh berbagai faktor lain, seperti komplikasi penyakit, respons terhadap terapi, maupun kondisi penyerta yang dimiliki pasien. Selain itu, kinerja algoritma KNN sangat bergantung pada penentuan nilai *k* atau jumlah tetangga terdekat yang digunakan dalam proses klasifikasi. Ketidaktepatan dalam memilih nilai *k*

dapat memicu kesalahan dalam pengelompokan kelas data, yang pada akhirnya berdampak pada penurunan akurasi model. Dengan demikian, pada penelitian ini KNN dinilai kurang memberikan hasil yang optimal jika dibandingkan dengan algoritma *Naive bayes* dan *Decision tree* (Tuloli et al., 2025).

5. PENUTUP

Berdasarkan hasil penelitian mengenai prediksi lama rawat inap pasien Diabetes Mellitus Tipe 2 menggunakan metode Knowledge Discovery in Databases (KDD) dengan algoritma *Naive bayes*, *Decision tree* dan *K-nearest neighbor* (KNN), dapat disimpulkan bahwa proses KDD berhasil digunakan dalam pengolahan data melalui tahapan *data selection*, *preprocessing*, transformasi data, data mining, serta evaluasi hasil model. Sehingga, dapat disimpulkan bahwa:

- a. Berdasarkan hasil pengujian algoritma *Naive bayes*, *Decision tree*, dan *K-nearest neighbor* (KNN) pada prediksi lama rawat inap pasien Diabetes Mellitus Tipe 2, dapat disimpulkan bahwa ketiga algoritma berhasil diterapkan dalam proses klasifikasi lama rawat inap pasien. Model mampu melakukan prediksi terhadap kategori lama rawat inap berdasarkan data pasien yang digunakan pada penelitian. Dari ketiga metode tersebut, algoritma *Naive bayes* menunjukkan keberhasilan penerapan terbaik karena mampu menghasilkan prediksi yang lebih stabil dan lebih baik dibandingkan algoritma lainnya. *Decision tree* juga mampu melakukan klasifikasi dengan cukup baik melalui pembentukan aturan pohon keputusan, sedangkan KNN masih kurang optimal karena karakteristik data medis yang kompleks dan distribusi data yang tidak seimbang.
- b. Berdasarkan hasil evaluasi yang dilakukan, algoritma *Naive bayes* menghasilkan nilai akurasi tertinggi yaitu sebesar 0,525 dibandingkan dengan *Decision tree* yang memperoleh 0,501 dan *K-nearest neighbor* (KNN) sebesar 0,438. Temuan ini menunjukkan bahwa *Naive bayes* memiliki kemampuan prediktif yang lebih baik dalam menentukan kategori lama rawat inap

pasien Diabetes Mellitus Tipe 2, karena mampu menghasilkan prediksi yang benar pada lebih dari separuh total data pengujian. Sementara itu, algoritma *Decision tree* menunjukkan kinerja yang cukup baik meskipun tingkat ketepatan prediksinya masih berada pada kategori sedang. Sementara itu, algoritma KNN memperoleh nilai accuracy paling rendah sehingga dinilai kurang optimal dalam memprediksi lama rawat inap pasien pada penelitian ini.

6. REFERENSI

- Azhar, Y., Firdausy, A. K., & Amelia, P. J. (2022). Perbandingan Algoritma Klasifikasi Data Mining Untuk Prediksi Penyakit Stroke. 5(2), 191–197.
- Bilqisth, S. C., Khoirudin, & Putri, A. N. (2022). Mengukur Tingkat Kepuasan Mahasiswa Terhadap *E- Learning* Universitas Semarang Menggunakan Algoritma *Naive Bayes*. 17, 1–7.
- Biyantoro, A. S., & Prasetyo, B. (2024). Penerapan *Decision tree* untuk Klasifikasi Status Kesehatan dengan perbandingan KNN dan *Naive bayes*. 4(1), 47–55.
- Fazriansyah, A., Alkhalifi, Y., & Zumarniansyah, A. (2025). Penerapan *Decision tree* Dengan Penyeimbangan Data *Imbalance* Menggunakan *Upsampling* Dalam Prediksi Penyakit Liver. 19(2), 259–266.
- Giusman, R., & Nurwahyuni, A. (2022). Biaya Pengobatan Pasien Rawat Inap Covid-19 Di Rumah Sakit X Tahun 2021. *Jurnal Ekonomi Kesehatan Indonesia*, 7(2), 96. <https://doi.org/10.7454/eki.v7i2.5797>
- Hunafa, M. R., & Hermawan, A. (2023). Perbandingan Algoritma *Naive Bayes* dan *K-nearest neighbor* Pada *Imbalance Class* Dataset Penyakit Diabetes. 4(3), 1551–1561. <https://doi.org/10.30865/klik.v4i3.1486>
- Kemendes RI. (2022). Permenkes No 24 Tahun 2022 Tentang Rekam Medis. 9, 356–363.
- Kementrian Kesehatan RI. (2022). Permenkes Ri No.24 Tahun 2022 Tentang Rekam Medis.
- Limbong, K., & Afriani, S. (2023). Hubungan Antara Waktu Tunggu Dengan Tingkat Kepuasan Pasien. *Flobamora Nursing Jurnal*, 3(1).

- Luxiarti, R., Aprilia, V. E., & Syaripudin, A. (2022). Analisis Indikator Lama Rawat Pasien Pada Kasus Typhoid Di Rumah Sakit Umum Daerah Arjawinangun Kabupaten Cirebon Analysis. *10*(2), 99–104.
- Nasien, D., Darwin, R., Cia, A., Winata, A. L., Go, J., M.C, R., Wijaya, R. C., & Lo, K. C. (2024). Perbandingan Implementasi *Machine Learning* Menggunakan Metode KNN , *Naive bayes* , Dan *Logistik Regression* Untuk Mengklasifikasi Penyakit Diabetes. *4*(1).
- Nugroho, B. I., Ma'arif, Z., & Arif, Z. (2022). Tinjauan Pustaka Sistematis: Penerapan Data Mining Metode Klasifikasi Untuk Menganalisa Penyalahgunaan Sosial Media. *Jurnal Sistem Informasi Dan Teknologi Peradaban*, *3*(2), 46–51.
- Rahayu, C. A. (2023). Prediksi Penderita Diabetes Menggunakan Metode *Naive bayes*. *Jurnal Informatika Dan Teknik Elektro Terapan*, *11*(3). <https://doi.org/10.23960/jitet.v11i3.3055>
- Ramadhani, A. A. A., Asriati, & Sety, L. O. M. (2025). Hubungan Rasionalitas Penggunaan Antibiotik terhadap Lama Rawat Inap pada Pasien Pneumonia Dewasa di RSUD Kota Kendari *Association Between Rational Antibiotic Use and Length of Stay in Adult Patients with Pneumonia at Kendari City Regional Hospital*. *17*(3), 267–276.
- Romlah, S. N., & Laiman, I. (2022). Hubungan Tingkat Pengetahuan Terhadap Perilaku Swamedikasi Obat Analgesik , Masyarakat Rw 04 Desa Trembulrejo Blora Periode April Tahun 2021. *IV*(1), 30–39.
- Sanusi, A., Putra, C. A., & Akbar, F. A. (2023). *Implementation Of Adaboost Algorithm On C50 For Improving The Performance Of Liver Disease Classification*. *8*(2), 93–102.
- Syahputra, P. (2022). Prediksi Lama Rawat Pasien Covid-19 Berbasis Machine Learning. *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, *9*(4), 3374–3382. <https://doi.org/10.35957/jatisi.v9i4.2883>
- Tuloli, M. S., Kinanti, T. S., & Amali, L. N. (2025). Perbandingan algoritma *C4 . 5* , *Naïve Bayes* , dan *K-nearest neighbors* untuk prediksi penyakit jantung. *7*(1), 11–21. <https://doi.org/10.37905/jji.v1i1.31158>
- Widodo, A. B., & Swedia, E. R. (2024). Analisis Perbandingan Algoritma KNN dan *Naïve Bayes* dalam Mendiagnosis Penyakit Diabetes Mellitus Pendahuluan. *23*(September), 387–396.
- Yunardi, R. T., & Dina, N. zata. (2022). *Data Mining dan Machine Learning dengan Orange3*. Airlangga University Press.

