



Artikel

Sistem Pendukung Keputusan untuk Klasifikasi Risiko Diabetes Menggunakan Algoritma *Decision Tree*

Hafizah Zuriyat Tayyibah^{1*}, Toat Tuloh¹, Khoirun Nisa¹, Glagah Eskacakra Setyowisnu², Rosyid Ridlo Al-Hakim³, Esa Rinjani Cantika Putri⁴

¹Program Studi Informatika, Universitas Harapan Bangsa, Purwokerto, Indonesia

²Program Studi Matematika, Universitas Jenderal Soedirman, Purwokerto, Indonesia

³Program Studi Sistem Informasi, Universitas Harapan Bangsa, Purwokerto, Indonesia

⁴Program Studi Kearsipan, Universitas Terbuka, Tangerang, Indonesia

* Korespondensi: hafizah.zuriyat@student.uhb.ac.id

Received: 1 Juli 2025

Revised: 30 Juli 2022

Accepted: 17 Agustus 2025

Published: 19 Agustus 2025



Copyright: © 2025 by the authors.

License Universitas Harapan Bangsa, Purwokerto, Indonesia. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license.

Abstrak: Diabetes melitus merupakan penyakit metabolik kronis yang memerlukan diagnosis dini untuk mencegah komplikasi serius. Penelitian ini bertujuan untuk membangun model klasifikasi risiko diabetes menggunakan algoritma *Decision Tree* berdasarkan dataset dari *Kaggle* yang terdiri dari 520 data pasien dengan 17 fitur klinis dan demografis, yang mana dapat dimanfaatkan oleh dokter spesialis endokrinologi dalam melakukan pemeriksaan. Hasil evaluasi menunjukkan bahwa model yang dikembangkan memiliki performa tinggi. Capaian ini menjadikan model lebih unggul dibanding enam penelitian terdahulu yang menggunakan algoritma seperti C4.5, *SVM*, *Adaboost*, *KNN*, dan *Naïve Bayes*, baik dari segi akurasi maupun sensitivitas deteksi terhadap kasus diabetes. Selain hasil metrik yang tinggi, penelitian ini juga menonjol dalam aspek interpretabilitas melalui visualisasi pohon keputusan yang mempermudah pemahaman logika klasifikasi oleh tenaga medis. Dengan cakupan atribut yang lebih luas dan evaluasi performa yang komprehensif, model ini dinilai efektif dan layak digunakan sebagai dasar pengembangan sistem pendukung keputusan untuk diagnosis dini diabetes yang akurat dan transparan.

Kata Kunci: sistem pendukung keputusan; *decision tree*; diabetes melitus; klasifikasi

Pendahuluan

Diabetes mellitus (DM) merupakan penyakit metabolik kronis yang prevalensinya meningkat secara global, dengan dampak yang signifikan terhadap kualitas hidup dan beban biaya kesehatan masyarakat (Camerlingo et al., 2022; Ratre et al., 2024). Penyakit ini disebabkan oleh gangguan pada sistem produksi atau respons tubuh terhadap insulin, yang mengakibatkan tingginya kadar glukosa dalam darah (Hardianto, 2021). Jika tidak ditangani dengan tepat, kondisi ini dapat menimbulkan komplikasi serius seperti penyakit jantung koroner, gagal ginjal, neuropati, dan kebutaan (Sriyati, 2024). Oleh karena itu, diagnosis dini sangat penting dilakukan untuk memungkinkan intervensi medis lebih awal guna mencegah dampak jangka panjang.

Seiring perkembangan teknologi, metode berbasis kecerdasan buatan seperti *machine learning* dan sistem pendukung keputusan (SPK) telah banyak dimanfaatkan dalam diagnosis penyakit, termasuk diabetes (Fadhillah et al., 2022). Salah satu algoritma yang populer dan banyak digunakan adalah *Decision Tree*, karena menghasilkan struktur model berbentuk pohon keputusan yang mudah dipahami dan diinterpretasikan, bahkan oleh tenaga medis non-teknis (Charbuty & Abdulazeez, 2021; Thakur et al., 2023).

Berbagai studi sebelumnya telah mengevaluasi efektivitas metode klasifikasi dalam mendeteksi diabetes. Misalnya, (Fadhillah et al., 2022) menerapkan algoritma C4.5 dan berhasil mencapai akurasi sebesar 76%, lebih tinggi dibanding algoritma SVM yang hanya mencapai 70% pada dataset *Pima Indians Diabetes*. Sementara itu, (Nurussakinah & Faisal, 2023) menggunakan algoritma *Decision Tree* dan memperoleh nilai *precision* sebesar 0.78, *recall* 0.45, dan *F1-score* 0.57, yang menunjukkan bahwa masih ada kelemahan model dalam mendeteksi kasus diabetes positif secara sensitif.

Di sisi lain, *Adaboost classifier* yang dikembangkan oleh (Abdurrahman, 2022) menunjukkan hasil yang cukup tinggi dengan akurasi 80.09% setelah proses imputasi data dengan metode *mean*, meskipun performanya sensitif terhadap nilai kosong (*missing values*). (Mucholladin et al., 2021) mengimplementasikan dua model *Support Vector Machine* (SVM), yaitu model *benchmark* dan model *scratch*, dengan akurasi masing-masing 87% dan 78%, serta *recall* model terbaik mencapai 78%. Selain itu, (Indrayanti et al., 2017) mengoptimalkan nilai K dalam algoritma *K-Nearest Neighbor* dan menemukan nilai optimal pada K=13 yang menghasilkan akurasi 75.14%, sedangkan (Ridwan, 2020) melaporkan bahwa *Naïve Bayes* mampu mencapai akurasi 90.20% dan *AUC* 0.95 dalam mendeteksi risiko diabetes pada fase awal.

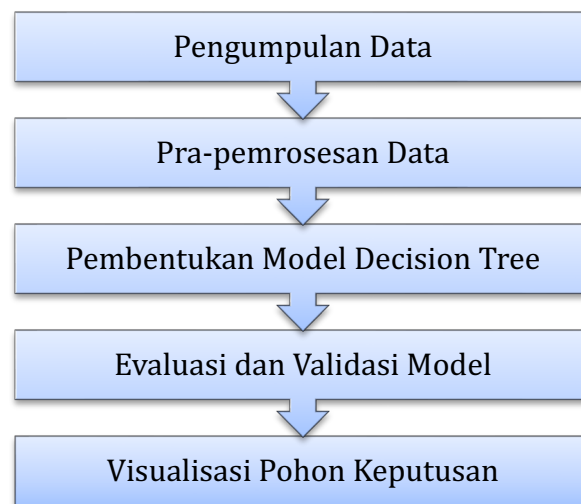
Melanjutkan hasil-hasil penelitian tersebut, studi ini bertujuan untuk mengembangkan model klasifikasi diabetes menggunakan algoritma *Decision Tree* dengan memperluas cakupan atribut, serta mengevaluasi kinerja model secara komprehensif menggunakan metrik akurasi, *precision*, *recall*, *f1-score*, dan *confusion matrix*. Dataset yang digunakan dalam penelitian ini berasal dari platform *Kaggle* dan memuat 520 data pasien dengan 17 fitur yang mencakup faktor demografis dan gejala klinis.

Selain menekankan pada akurasi model, penelitian ini juga menitikberatkan pada interpretabilitas hasil klasifikasi melalui visualisasi pohon keputusan. Visualisasi ini diharapkan dapat membantu tenaga medis untuk memahami logika pengambilan keputusan yang dihasilkan oleh model, sehingga dapat menjadi dasar pertimbangan dalam proses diagnosis. Dengan pendekatan yang lebih menyeluruh dan berbasis data global, penelitian ini diharapkan dapat memberikan kontribusi dalam pengembangan sistem pendukung keputusan yang akurat, transparan, dan aplikatif dalam lingkungan klinis untuk diagnosis dini penyakit diabetes. Pemilihan algoritma *Decision Tree* dalam penelitian ini didasarkan pada kelebihanannya dalam menangani fitur campuran

(numerik dan kategorikal), kemudahan interpretasi hasil, serta kemampuannya dalam mengidentifikasi aturan klasifikasi secara eksplisit tanpa memerlukan transformasi data kompleks. Dibandingkan algoritma lain seperti SVM atau *Neural Network*, *Decision Tree* lebih ringan secara komputasi dan cocok digunakan untuk aplikasi berbasis sistem pendukung keputusan yang memerlukan transparansi dan penelusuran logika keputusan.

Metode Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen berbasis *machine learning*, yaitu algoritma *Decision Tree*, untuk membangun sistem pendukung keputusan dalam klasifikasi risiko diabetes. Penelitian ini terdiri dari lima tahap utama sebagaimana ditunjukkan pada Gambar 1.



Gambar 1. Metode Penelitian

Pengumpulan Data

Data yang digunakan pada penelitian ini diperoleh dari dataset publik “*Diabetes Risk Prediction*” oleh Himanshu di platform *Kaggle* (<https://www.kaggle.com/datasets/rcratos/diabetes-risk-prediction>), yang berisi 520 entri data pasien. Masing-masing entri memuat 17 atribut, yang mencakup faktor demografis seperti usia dan jenis kelamin, serta gejala klinis seperti *polyuria*, *polydipsia*, *sudden weight loss*, dan *delayed healing*. Data ini bersifat anonim dan tidak mengandung informasi identitas pribadi, sehingga aman untuk digunakan sebagai bahan analisis (Sadiq et al., 2025). Keputusan untuk menggunakan dataset ini didasarkan pada kebutuhan akan data dengan variabel yang beragam dan representatif secara klinis. Hal ini cukup penting agar model dapat belajar dari variasi faktor risiko yang lebih luas, yang mana akan menghasilkan klasifikasi yang lebih akurat dan mendekati kondisi nyata.

Tabel 1. Atribut Dataset

<i>Atribut</i>	<i>Deskripsi</i>	<i>Nilai</i>
<i>Age</i>	Usia individu	Angka (misal: 25, 40)
<i>Gender</i>	Jenis kelamin	<i>Male, Female</i>

<i>Polyuria</i>	Sering buang air kecil	<i>Yes, No</i>
<i>Polydipsia</i>	Sering merasa haus	<i>Yes, No</i>
<i>Sudden weight loss</i>	Penurunan berat badan tiba-tiba	<i>Yes, No</i>
<i>Weakness</i>	Tubuh terasa lemah	<i>Yes, No</i>
<i>Polyphagia</i>	Lapar berlebihan	<i>Yes, No</i>
<i>Genital thrush</i>	Infeksi jamur di area genital	<i>Yes, No</i>
<i>Visual blurring</i>	Penglihatan kabur	<i>Yes, No</i>
<i>Itching</i>	Gatal-gatal	<i>Yes, No</i>
<i>Irritability</i>	Mudah marah / tersinggung	<i>Yes, No</i>
<i>Delayed healing</i>	Luka sulit sembuh	<i>Yes, No</i>
<i>Partial paresis</i>	Kelumpuhan sebagian	<i>Yes, No</i>
<i>Muscle stiffness</i>	Kekakuan otot	<i>Yes, No</i>
<i>Alopecia</i>	Rambut rontok	<i>Yes, No</i>
<i>Obesity</i>	Kegemukan / obesitas	<i>Yes, No</i>
<i>Class</i>	Kelas Positif dan Negatif	1, 0

Pra-pemrosesan Data

Sebelum dilakukan pelatihan model, data terlebih dahulu melalui tahap pra-pemrosesan guna memastikan bahwa seluruh variabel berada dalam format yang sesuai untuk diolah oleh algoritma klasifikasi. Tahap ini memastikan data siap diproses mesin tanpa kehilangan makna dari atribut. Salah satu tahap awal dalam pra-pemrosesan data ini adalah mengubah nilai pada kolom '*class*' dari format *string* ('*Negative*' dan '*Positive*') menjadi nilai *numerik* (0 dan 1) menggunakan fungsi `.map()`. Konversi tersebut diperlukan agar kompatibel dengan algoritma *sklearn*. Setelah konversi, dilakukan pengecekan distribusi kelas menggunakan fungsi `.value_counts()` untuk melihat apakah jumlah data pada masing-masing kelas relatif seimbang. Kedua langkah ini penting untuk memastikan integritas data sebelum masuk ke tahap pelatihan model. Proses lengkap dari konversi dan analisis distribusi kelas tersebut dapat dilihat pada Gambar 2.

```
# Ubah kolom 'class' menjadi 0 dan 1
df['class'] = df['class'].map({'Negative': 0, 'Positive': 1})

# Tampilkan distribusi kelas setelah dikonversi menjadi 0 dan 1
print("Distribusi Kelas (0 = Negative, 1 = Positive):")
print(df['class'].value_counts())
```

Gambar 2. Pengubahan Label Target

Setelah memastikan bahwa label target telah dikonversi ke dalam format numerik dan distribusinya seimbang, tahap selanjutnya adalah melakukan *encoding* pada fitur-fitur kategorikal. Semua fitur bertipe kategorikal seperti *Polyuria*, *Polydipsia*, *Itching*, dan *Polyphagia* diubah ke dalam bentuk numerik menggunakan teknik *one-hot encoding*, sebagaimana juga diterapkan oleh (Charbuty & Abdulazeez, 2021). Teknik ini digunakan agar nilai-nilai kategorikal dapat direpresentasikan secara numerik tanpa mengasumsikan urutan atau bobot antar kategori, sehingga lebih sesuai untuk algoritma seperti *Decision Tree*. Hasil dari proses *one-hot encoding* ditunjukkan pada Gambar 3.

```
Index(['Age', 'class', 'Gender_Male', 'Polyuria_Yes', 'Polydipsia_Yes',
      'sudden weight loss_Yes', 'weakness_Yes', 'Polyphagia_Yes',
      'Genital thrush_Yes', 'visual blurring_Yes', 'Itching_Yes',
      'Irritability_Yes', 'delayed healing_Yes', 'partial paresis_Yes',
      'muscle stiffness_Yes', 'Alopecia_Yes', 'Obesity_Yes'],
      dtype='object')
```

Gambar 3. Hasil *Encoding* pada Fitur Kategorikal

Langkah selanjutnya setelah proses *one-hot encoding* selesai dan seluruh fitur berada dalam format numerik adalah membagi dataset menjadi dua bagian, yaitu data latih dan data uji. Pembagian dilakukan menggunakan fungsi *train_test_split* dari pustaka *Scikit-learn* dengan rasio 80:20 dan parameter *random_state=42* untuk memastikan hasil yang konsisten di setiap eksekusi. Dari total 520 data pasien, sebanyak 416 data digunakan sebagai data latih, sedangkan 104 data sisanya digunakan sebagai data uji. Pemisahan ini penting untuk memastikan bahwa model hanya belajar dari data latih dan dievaluasi menggunakan data yang belum pernah dilihat sebelumnya, sehingga performa model dapat diukur secara objektif dan risiko *overfitting* dapat diminimalkan.

Pembentukan Model *Decision Tree*

Model klasifikasi dibentuk menggunakan algoritma *Decision Tree Classifier* dari pustaka *sklearn.tree*. Model ini dipilih karena kemampuannya dalam menghasilkan aturan klasifikasi yang mudah dipahami dan efektif untuk data dengan tipe campuran, baik numerik maupun kategorikal (Charbuty & Abdulazeez, 2021; Nurussakinah & Faisal, 2023). Dalam implementasi ini, model dibuat menggunakan parameter *default* dengan nilai *random_state=42* untuk memastikan reproduktibilitas hasil pelatihan. Pemilihan fitur terbaik pada setiap *node* dalam pohon dilakukan secara otomatis berdasarkan kriteria pemisahan *Gini Impurity*, yang merupakan metode *default* di *Scikit-learn*. Kriteria ini juga relatif lebih efisien secara komputasi, terutama jika diterapkan pada dataset dengan 17 atribut. Rumus untuk menghitung *Gini Impurity* pada suatu *node* dijelaskan pada Persamaan (1).

$$Gini(t) = 1 - \sum_{i=1}^c p_i^2 \quad (1)$$

Keterangan:

- $Gini(t)$: nilai *impurity* pada *node t*
- c : jumlah kelas
- p_i : proporsi data dari kelas ke- i pada *node* tersebut

Semakin kecil nilai *Gini*, semakin baik suatu atribut dalam memisahkan data ke dalam kelas-kelas target. Proses pelatihan dilakukan pada data latih menggunakan perintah *model.fit(X_train, y_train)*, sedangkan hasil prediksi model kemudian dievaluasi menggunakan data uji.

Evaluasi dan Validasi Model

Evaluasi performa model dilakukan untuk mengukur sejauh mana algoritma *Decision Tree* mampu mengklasifikasikan risiko diabetes dengan benar terhadap data uji. Proses evaluasi dimulai dengan melakukan prediksi menggunakan *model.predict(X_test)*, yang hasilnya disimpan dalam variabel *y_pred*. Selanjutnya, hasil prediksi dibandingkan dengan data aktual (*y_test*) menggunakan fungsi

`accuracy_score`, `classification_report`, dan `confusion_matrix` dari pustaka *Scikit-learn*. Evaluasi ini menghasilkan empat metrik utama, yaitu akurasi, *precision*, *recall*, dan *f1-score*, yang masing-masing menggambarkan aspek berbeda dari performa model.

- a. Akurasi mengukur proporsi prediksi yang benar dari seluruh data, dirumuskan pada Persamaan (2).

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

- b. *Precision* mengukur ketepatan model dalam mengklasifikasikan kelas positif (Persamaan (3)).

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

- c. *Recall* mengukur seberapa banyak data positif yang berhasil terdeteksi (Persamaan (4)).

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

- d. *F1-score* merupakan harmonisasi antara *precision* dan *recall* (Persamaan (5)).

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

Keterangan:

- TP (*True Positive*) : jumlah kasus positif yang diprediksi benar;
- TN (*True Negative*) : jumlah kasus negatif yang diprediksi benar;
- FP (*False Positive*) : jumlah kasus negatif yang salah diprediksi sebagai positif oleh model;
- FN (*False Negative*) : jumlah kasus positif yang salah diprediksi sebagai negatif oleh model.

Keempat metrik ini ditampilkan secara otomatis melalui `classification_report(y_test, y_pred)` dalam format tabel. Selain itu, untuk visualisasi hasil klasifikasi, `confusion matrix` divisualisasikan menggunakan fungsi `heatmap()` dari pustaka *Seaborn*. Melalui `confusion matrix` ini, pengguna dapat dengan jelas melihat jumlah prediksi benar dan salah untuk masing-masing kelas, sehingga evaluasi model tidak hanya bersifat numerik, tetapi juga visual dan interpretatif.

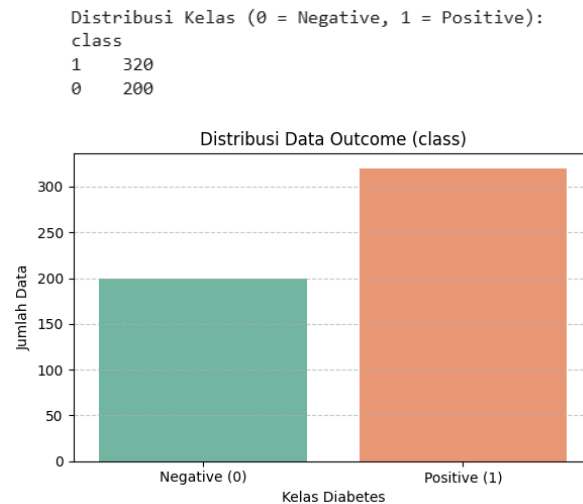
Visualisasi Pohon Keputusan

Langkah terakhir dari metode ini adalah memvisualisasikan struktur pohon keputusan yang telah dilatih menggunakan fungsi `plot_tree()` dari pustaka *Scikit-learn*. Visualisasi ini bertujuan untuk memberikan interpretasi terhadap alur logika klasifikasi yang dibentuk oleh model, sehingga hasil klasifikasi tidak hanya menjadi keluaran numerik, tetapi juga dapat dipahami oleh praktisi kesehatan yang memerlukan transparansi dalam proses pengambilan keputusan. Pendekatan visual seperti ini telah digunakan dalam studi oleh (Fadhilah et al., 2022), yang menunjukkan bahwa model yang bersifat *interpretable* dapat meningkatkan kepercayaan pengguna terhadap sistem berbasis kecerdasan buatan, terutama dalam aplikasi medis. Lebih lanjut, pengguna dapat langsung menghubungkan hasil model dengan gejala klinis nyata.

Hasil dan Pembahasan

Pra-pemrosesan Data

Berikut adalah *output* dari perubahan nilai pada kolom '*class*' dari format *string* ('*Negative*' dan '*Positive*') menjadi nilai numerik (0 dan 1) beserta visualisasinya dalam bentuk diagram.



Gambar 4. Visualisasi Distribusi Kelas

Visualisasi distribusi kelas setelah konversi label target ke format numerik ditampilkan pada Gambar 4. Diagram batang ini menunjukkan jumlah data pasien berdasarkan dua kategori utama, yaitu kelas *Negative* (0) dan *Positive* (1), yang masing-masing merepresentasikan pasien non-diabetes dan pasien dengan risiko diabetes. Dari gambar tersebut terlihat bahwa jumlah data pada kelas positif lebih tinggi, yaitu sekitar 320 pasien, sedangkan kelas negatif berjumlah sekitar 200 pasien.

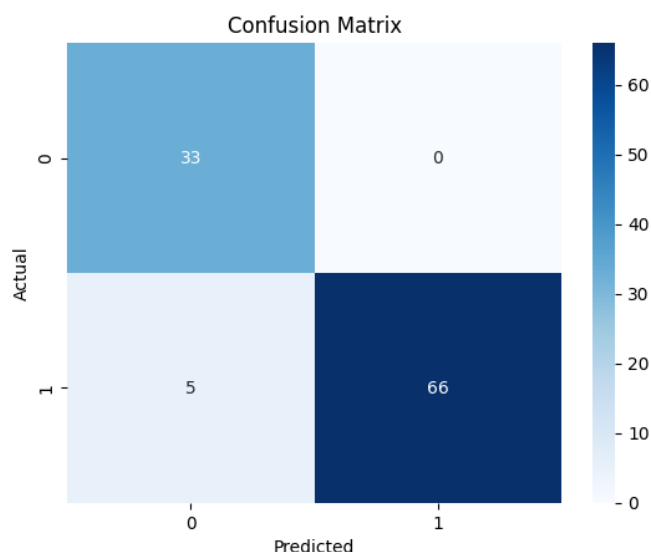
Pelatihan dan Evaluasi Model

Setelah dilakukan pelatihan dan evaluasi model *Decision Tree* pada data uji, diperoleh hasil klasifikasi yang menunjukkan performa cukup tinggi. Model ini mampu melakukan klasifikasi risiko diabetes secara akurat dengan mempertimbangkan kombinasi dari berbagai fitur demografis dan gejala klinis. Evaluasi model dalam penelitian ini dilakukan menggunakan empat metrik utama, yaitu akurasi, *precision*, *recall*, dan *f1-score*. Seluruh metrik dihitung berdasarkan data uji sebanyak 104 data pasien. Hasil evaluasi lengkap disajikan dalam Tabel 2 berikut.

Tabel 2. Hasil Evaluasi Model *Decision Tree*

Kelas	Precision	Recall	F1-Score	Support
Negatif (0)	0.87	1.00	0.93	33
Positif (1)	1.00	0.93	0.96	71
Accuracy	-	-	0.95	104
Macro Avg	0.93	0.96	0.95	104
Weighted Avg	0.96	0.95	0.95	104

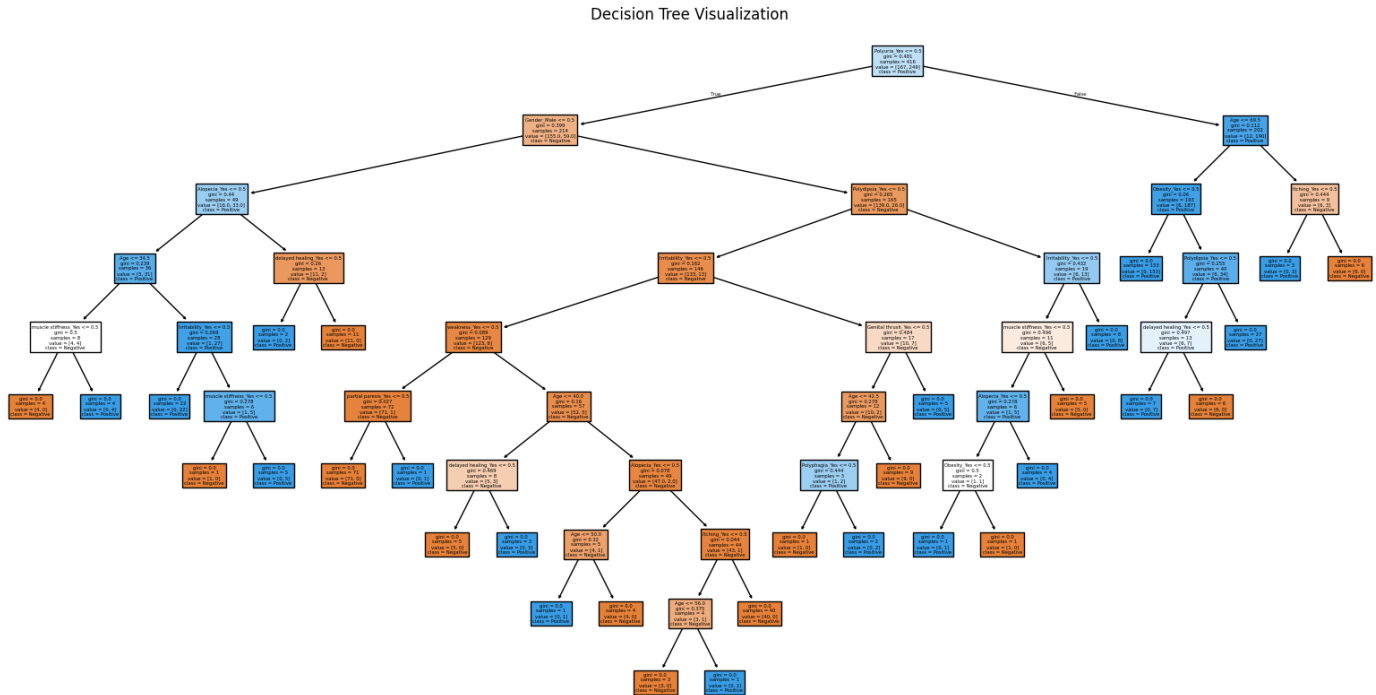
Hasil pada Tabel 2 menunjukkan bahwa model memiliki akurasi keseluruhan sebesar 95,19% yang mana lebih tinggi dari penelitian-penelitian sebelumnya. Kemudian diperoleh pula *precision* 1.00 pada kelas positif, yang berarti tidak ada kasus negatif yang salah diklasifikasikan sebagai positif (*false positive* = 0). Hal ini memberikan dampak di mana kesalahan diagnosis pada pasien sehat dapat dihindari. Kemudian *recall* untuk kelas positif mencapai 0.93, yang berarti model berhasil mendeteksi sebagian besar pasien dengan diabetes secara akurat. Lebih lanjut, nilai *f1-score* sebesar 0.96 menunjukkan model yang digunakan sangat baik dalam klasifikasi, di mana hampir setiap prediksi yang dilakukan positif benar. Hal ini menunjukkan bahwa model ini cocok diaplikasikan dalam bidang medis, khususnya diabetes.



Gambar 5. Visualisasi *Confusion Matrix*

Visualisasi dari *Confusion Matrix* pada Gambar 5 menunjukkan bahwa dari 71 data pasien yang benar-benar positif diabetes, sebanyak 66 berhasil diklasifikasikan dengan benar dan hanya 5 yang terklasifikasi salah sebagai negatif. Ini sangat penting dalam konteks medis, di mana kesalahan dalam mendeteksi pasien positif (*false negative*) dapat berakibat pada keterlambatan diagnosis dan penanganan klinis.

Visualisasi Pohon Keputusan



Gambar 6. Visualisasi Pohon Keputusan

Hasil visualisasi pohon keputusan yang ditampilkan pada Gambar 6 memperlihatkan struktur model dalam bentuk hierarki cabang yang menggambarkan proses pemisahan data berdasarkan fitur-fitur penting. Pada bagian akar pohon, terlihat bahwa fitur seperti *Polyuria_Yes* dan *Polydipsia_Yes* menjadi pemisah utama, yang menandakan bahwa gejala sering buang air kecil dan rasa haus berlebih merupakan indikator awal yang sangat kuat dalam klasifikasi risiko diabetes. Cabang-cabang selanjutnya menunjukkan kombinasi fitur-fitur lain seperti *Itching*, *Delayed healing*, *Age*, dan *Gender*, yang juga memainkan peran penting dalam keputusan klasifikasi.

Warna pada *node* digunakan untuk menunjukkan kelas mayoritas yang diprediksi pada setiap simpul. Warna oranye mewakili kelas positif (penderita diabetes) dan warna biru mewakili kelas negatif (non-diabetes). Semakin gelap warnanya, semakin kuat dominasi kelas tersebut dalam *node* tersebut. Setiap simpul juga dilengkapi dengan informasi seperti jumlah sampel (*samples*), nilai *Gini Impurity* (*gini*), dan distribusi kelas (*value*), yang dapat digunakan untuk menganalisis kontribusi masing-masing atribut secara lebih mendalam. Dengan demikian, visualisasi ini tidak hanya meningkatkan interpretabilitas model, tetapi juga memberikan wawasan tambahan bagi praktisi medis dalam memahami korelasi antar gejala dan hasil diagnosis. Namun demikian, model ini memiliki beberapa keterbatasan, seperti potensi *overfitting* apabila pohon terlalu dalam, serta ketergantungan terhadap distribusi data yang digunakan. Distribusi data yang tidak seimbang antara kelas positif dan negatif juga dapat mempengaruhi sensitivitas model. Oleh karena itu, pengembangan lebih lanjut perlu mempertimbangkan teknik *pruning* dan peningkatan generalisasi model agar hasil klasifikasi tetap andal pada data populasi berbeda.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa algoritma *Decision Tree* mampu menghasilkan klasifikasi yang akurat, sensitif, dan mudah diinterpretasikan, menjadikannya cocok untuk digunakan dalam sistem pendukung keputusan diagnosis dini diabetes. Kelebihan lainnya adalah visualisasi pohon yang jelas, sehingga model ini dapat dijadikan alat bantu analisis oleh tenaga medis tanpa kehilangan transparansi logika keputusan. Dengan performa yang lebih baik dibandingkan penelitian sebelumnya, model ini berpotensi untuk diintegrasikan dalam sistem klinis sebagai bagian dari alat skrining atau edukasi pasien berbasis data.

Kesimpulan

Berdasarkan hasil evaluasi model, penelitian ini menunjukkan performa yang sangat baik dalam mengklasifikasikan risiko diabetes menggunakan algoritma *Decision Tree*. Model yang dibangun berhasil mencapai akurasi sebesar 95,19%, dengan *precision* 1.00, *recall* 0.93, dan *f1-score* 0.96 pada kelas positif (penderita diabetes). Dibandingkan dengan enam penelitian terdahulu, model ini memiliki keunggulan yang signifikan. (Fadhillah et al., 2022) melaporkan akurasi 76% dengan C4.5, sementara (Nurussakinah & Faisal, 2023) memperoleh *precision* 0.78 dan *recall* yang masih rendah sebesar 0.45. Metode *Adaboost* oleh (Abdurrahman, 2022) menghasilkan akurasi 80.09%, dan SVM oleh (Mucholladin et al., 2021) mencapai akurasi 87% dengan *recall* 0.78. Selain itu, model KNN oleh (Indrayanti et al., 2017) dan *Naïve Bayes* oleh (Ridwan, 2020) masing-masing memperoleh akurasi 75.14% dan 90.20%.

Keunggulan utama dari penelitian ini terletak pada kombinasi akurasi tinggi, keseimbangan metrik *precision* dan *recall*, serta interpretabilitas model melalui visualisasi pohon keputusan. Selain itu, penelitian ini menggunakan cakupan fitur yang lebih luas (17 atribut) dibanding studi terdahulu yang umumnya hanya menggunakan 8–9 atribut, sehingga meningkatkan kekayaan informasi dalam proses klasifikasi. Dengan demikian, model yang dikembangkan tidak hanya unggul secara kuantitatif, tetapi juga mendukung transparansi dan kemudahan interpretasi bagi tenaga medis. Penelitian ini diharapkan dapat berkontribusi pada pengembangan sistem pendukung keputusan yang efektif dan akurat untuk diagnosis dini penyakit diabetes. Langkah selanjutnya yang disarankan dari penelitian ini adalah pengembangan sistem pendukung keputusan (SPK) berbasis web atau desktop yang mengintegrasikan model klasifikasi ini ke dalam antarmuka pengguna interaktif. Hal ini akan memungkinkan tenaga medis atau pasien untuk melakukan deteksi awal secara mandiri dengan interpretasi hasil yang mudah dipahami. Selain itu, pengujian sistem pada data klinis nyata secara prospektif perlu dilakukan untuk validasi eksternal sebelum implementasi secara luas.

Referensi

- Abdurrahman, G. (2022). Jurnal Sistem dan Teknologi Informasi Klasifikasi Penyakit Diabetes Melitus Menggunakan Adaboost Classifier. *JUSTINDO (Jurnal Sistem Dan Teknologi Informasi)*, 7(1), 59–66. <http://jurnal.unmuhjember.ac.id/index.php/JUSTINDO>
- Camerlingo, N., Vettoretti, M., Del Favero, S., Facchinetti, A., Choudhary, P., & Sparacino, G. (2022). Generation of Post-Meal Insulin Correction Boluses in Type 1 Diabetes Simulation Models for In-Silico Clinical Trials: More Realistic Scenarios Obtained using a Decision Tree Approach. *Computer Methods and Programs in Biomedicine*, 221, 106862. <https://doi.org/10.1016/j.cmpb.2022.106862>
- Charbuty, B., & Abdulazeez, A. (2021). Classification Based on Decision Tree Algorithm for Machine Learning. *Journal of Applied Science and Technology Trends*, 2(01), 20–28. <https://doi.org/10.38094/jastt20165>
- Fadhillah, R. P., Rahma, R., Sepharni, A., Mufidah, R., Sari, B. N., & Pangestu, A. (2022). Klasifikasi Penyakit Diabetes

- Mellitus Berdasarkan Faktor-Faktor Penyebab Diabetes menggunakan Algoritma C4.5. *JUPI (Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika)*, 7(4), 1265–1270. <https://doi.org/10.29100/jipi.v7i4.3248>
- Hardianto, D. (2021). Telaah Komprehensif Diabetes Melitus: Klasifikasi, Gejala, Diagnosis, Pencegahan, dan Pengobatan. *Jurnal Bioteknologi & Biosains Indonesia (JBBi)*, 7(2), 304–317. <https://doi.org/10.29122/jbbi.v7i2.4209>
- Indrayanti, Devi Sugianti, & M. Adib Al Karomi. (2017). Optimasi Parameter K Pada Algoritma K-Nearest Neighbour untuk Klasifikasi Penyakit Diabetes Mellitus. *Prosiding Seminar Nasional Teknologi Dan Informatika*, 4, 823–830.
- Mucholladin, A. W., Bachtiar, F. A., & Furqon, M. T. (2021). Klasifikasi Penyakit Diabetes menggunakan Metode Support Vector Machine. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 5(2), 622–633. <http://j-ptiik.ub.ac.id>
- Nurussakinah, N., & Faisal, M. (2023). Klasifikasi Penyakit Diabetes Menggunakan Algoritma Decision Tree. *Jurnal Informatika*, 10(2), 143–149. <https://doi.org/10.31294/inf.v10i2.15989>
- Ratre, B., Patel, D., Patel, D., Patel, H., Patel, S., & Singh, D. (2024). Review Paper for Diabetes Mellitus. *International Journal of Advanced Multidisciplinary Research and Studies*, 4(6), 927–931. <https://doi.org/10.62225/2583049X.2024.4.6.3534>
- Ridwan, A. (2020). Penerapan Algoritma Naïve Bayes untuk Klasifikasi Penyakit Diabetes Mellitus. *Jurnal SISKOM-KB (Sistem Komputer Dan Kecerdasan Buatan)*, 4(1), 15–21. <https://doi.org/10.47970/siskom-kb.v4i1.169>
- Sadiq, I. Z., Katsayal, B. S., Ibrahim, B., Ibrahim, M., Hassan, H. A., Ghali, U. M., Usman, A. G., Usman, A., & Abba, S. I. (2025). Data-Driven Diabetes Mellitus Prediction and Management: A Comparative Evaluation of Decision Tree Classifier and Artificial Neural Network Models Along with Statistical Analysis. *Scientific Reports*, 15(1), 19339. <https://doi.org/10.1038/s41598-025-03718-w>
- Sriyati, S. (2024). Neuropati Diabetes Sebagai Faktor Predisposisi Terjadinya Luka Pada Kaki. *Jurnal Ilmiah STIKES Yarsi Mataram*, 14(1), 46–52. <https://doi.org/10.57267/jisym.v14i1.336>
- Thakur, R., . M., Sharma, P., & . D. (2023). Diseases Prediction Using Classification Algorithm. *International Journal for Research in Applied Science and Engineering Technology*, 11(8). <https://doi.org/10.22214/ijraset.2023.55194>