



Implementasi dan Evaluasi Performa Algoritma Naïve Bayes dalam Deteksi Dini Penyakit Diabetes

Nurhasanah*, Nilovar Asyiah, Rahmawati

Fakultas Ilmu Komputer, Program Studi Teknik Informatika, Universitas Pamulang, Tangerang Selatan, Indonesia

Email: ^{1,*}dosen02834@unpam.ac.id, ²dosen02835@unpam.ac.id, ³dosen02394@unpam.ac.id

Email Penulis Korespondensi: dosen02834@unpam.ac.id

Abstrak—Diabetes mellitus merupakan salah satu penyakit kronis yang menjadi penyebab utama meningkatnya angka morbiditas dan mortalitas di berbagai negara. Deteksi dini terhadap risiko diabetes sangat penting untuk membantu pencegahan komplikasi melalui penanganan yang lebih cepat dan tepat. Penelitian ini bertujuan mengimplementasikan dan mengevaluasi performa algoritma Naïve Bayes dalam mendukung deteksi dini penyakit diabetes berdasarkan data kesehatan pasien. Dataset yang digunakan adalah Pima Indians Diabetes Dataset yang terdiri atas 768 data dengan delapan atribut masukan dan satu atribut keluaran. Tahap preprocessing dilakukan melalui penanganan nilai nol (zero values) pada atribut yang secara fisiologis tidak mungkin bernilai nol menggunakan pendekatan median, dilanjutkan dengan pemeriksaan konsistensi data dan pembagian data menggunakan metode split validation dengan rasio 80:20. Evaluasi model dilakukan menggunakan confusion matrix dengan parameter accuracy, precision, recall, dan F1-score. Hasil penelitian menunjukkan bahwa algoritma Naïve Bayes memperoleh nilai accuracy sebesar 88,31%, precision sebesar 87,80%, recall sebesar 90,00%, dan F1-score sebesar 88,89%. Hasil tersebut menunjukkan bahwa algoritma Naïve Bayes memiliki performa yang baik dalam mengklasifikasikan risiko diabetes. Implementasi model ke dalam aplikasi berbasis web diharapkan dapat dimanfaatkan sebagai alat bantu skrining awal (early screening) oleh tenaga kesehatan maupun masyarakat untuk mendukung pengambilan keputusan awal sebelum dilakukan pemeriksaan medis lebih lanjut.

Kata Kunci: Diabetes Mellitus; Naïve Bayes; Klasifikasi; Machine Learning; Deteksi Dini

Abstract—Diabetes mellitus is one of the most prevalent chronic diseases worldwide and requires early detection to reduce the risk of severe complications through timely intervention. This study aims to implement and evaluate the performance of the Naïve Bayes algorithm in supporting the early detection of diabetes based on patients' health data. The study employed the Pima Indians Diabetes Dataset, consisting of 768 patient records with eight input attributes and one output attribute. During the preprocessing stage, zero values in physiological attributes were treated as missing values and replaced using the median of each respective attribute, followed by data consistency checking and dataset partitioning using the 80:20 split validation method. Model performance was evaluated using a confusion matrix with four performance metrics: accuracy, precision, recall, and F1-score. The experimental results showed that the Naïve Bayes algorithm achieved an accuracy of 88.31%, precision of 87.80%, recall of 90.00%, and an F1-score of 88.89%. These findings indicate that the proposed model performs well in classifying diabetes risk. The implementation of the model in a web-based application is expected to assist healthcare professionals and the general public as an early screening tool to support preliminary decision-making before comprehensive medical examination.

Keywords: Diabetes Mellitus; Naïve Bayes; Classification; Machine Learning; Early Detection

1. PENDAHULUAN

Diabetes mellitus merupakan salah satu penyakit kronis yang menjadi tantangan utama dalam bidang kesehatan karena prevalensinya terus meningkat di berbagai negara. Penyakit ini ditandai dengan gangguan metabolisme glukosa akibat ketidakmampuan tubuh memproduksi insulin dalam jumlah yang cukup atau menggunakan insulin secara efektif sehingga menyebabkan hiperglikemia. Apabila tidak ditangani sejak dini, diabetes dapat menimbulkan berbagai komplikasi serius, seperti penyakit kardiovaskular, gagal ginjal, neuropati, retinopati, hingga peningkatan risiko kematian. Menurut International Diabetes Federation (IDF), jumlah penyandang diabetes di dunia diperkirakan terus meningkat pada beberapa dekade mendatang sehingga diperlukan strategi deteksi dini yang lebih efektif untuk mengurangi risiko komplikasi dan meningkatkan kualitas hidup pasien [1]. Selain itu, American Diabetes Association (ADA) menegaskan bahwa identifikasi risiko diabetes sejak tahap awal merupakan bagian penting dalam pencegahan, pengelolaan penyakit, serta pengambilan keputusan klinis yang lebih tepat sehingga dapat menekan angka morbiditas dan mortalitas akibat diabetes [2].

Perkembangan *machine learning* telah memberikan kontribusi yang signifikan dalam pengembangan sistem prediksi penyakit diabetes melalui pemanfaatan data kesehatan pasien. Berbagai algoritma klasifikasi mampu mengenali pola pada data klinis sehingga proses identifikasi risiko diabetes dapat dilakukan secara lebih cepat, objektif, dan konsisten dibandingkan pendekatan konvensional. Tinjauan literatur menunjukkan bahwa penelitian mengenai penerapan *machine learning* dan *artificial intelligence* pada prediksi diabetes mengalami perkembangan yang sangat pesat dengan fokus pada pengembangan model yang memiliki performa klasifikasi tinggi [3]. Kajian lain juga menunjukkan bahwa pemanfaatan algoritma *machine learning* mampu meningkatkan efisiensi analisis data kesehatan serta mendukung pengambilan keputusan berbasis data pada proses deteksi dini diabetes [4]. Selain itu, penerapan algoritma *supervised learning* pada data rekam kesehatan elektronik menunjukkan performa yang baik dalam memprediksi risiko diabetes sehingga berpotensi dimanfaatkan sebagai pendukung proses skrining awal sebelum dilakukan pemeriksaan klinis lebih lanjut [5]. Oleh karena itu, pemilihan algoritma klasifikasi tidak hanya mempertimbangkan tingkat akurasi, tetapi juga efisiensi komputasi, kemudahan implementasi, serta kesesuaiannya terhadap karakteristik data yang digunakan.



Di antara berbagai algoritma klasifikasi yang digunakan dalam prediksi penyakit diabetes, Naïve Bayes merupakan salah satu metode yang banyak diterapkan karena memiliki proses komputasi yang sederhana, waktu pelatihan yang relatif cepat, serta mampu menghasilkan performa klasifikasi yang kompetitif pada berbagai kasus klasifikasi kesehatan. Pada data yang didominasi oleh atribut numerik kontinu, Gaussian Naïve Bayes menjadi salah satu varian yang sesuai karena memodelkan distribusi setiap atribut menggunakan distribusi Gaussian. Penelitian yang dilakukan oleh Resti et al. menunjukkan bahwa algoritma Naïve Bayes mampu digunakan untuk diagnosis diabetes mellitus dengan performa yang baik [6]. Hasil serupa juga dilaporkan oleh Salsabila et al. yang menunjukkan bahwa Naïve Bayes mampu menghasilkan klasifikasi diabetes dengan tingkat akurasi yang memadai [7]. Selain itu, Amal et al. menyimpulkan bahwa pendekatan probabilistik yang digunakan oleh Naïve Bayes tetap efektif dalam mengidentifikasi risiko diabetes berdasarkan data kesehatan pasien [8]. Temuan-temuan tersebut menunjukkan bahwa Naïve Bayes masih menjadi salah satu metode klasifikasi yang relevan untuk diterapkan pada prediksi penyakit diabetes.

Meskipun demikian, perkembangan penelitian dalam beberapa tahun terakhir lebih banyak berfokus pada penggunaan algoritma yang lebih kompleks untuk meningkatkan performa klasifikasi. Islam et al. melaporkan bahwa kombinasi *stacked machine learning* dan *deep learning* mampu meningkatkan akurasi prediksi diabetes [9]. Nguyen et al. menunjukkan bahwa pendekatan *hybrid deep learning* dapat meningkatkan kemampuan model dalam mengidentifikasi risiko prediabetes [10], sedangkan Majyambere et al. mengembangkan pendekatan *explainable machine learning* untuk meningkatkan interpretasi hasil prediksi sehingga model menjadi lebih mudah dipahami oleh pengguna [11]. Di sisi lain, Alzboon et al. menjelaskan bahwa pemilihan algoritma klasifikasi tidak hanya ditentukan oleh tingkat akurasi, tetapi juga harus mempertimbangkan efisiensi komputasi dan karakteristik data yang digunakan [12]. Hasil penelitian Khafaga et al. juga menunjukkan bahwa metode klasifikasi dengan kompleksitas yang lebih rendah masih mampu memberikan performa yang baik apabila diterapkan pada data kesehatan yang sesuai [13]. Selain itu, Mousa et al. menyatakan bahwa model prediksi yang sederhana lebih mudah diintegrasikan ke dalam aplikasi kesehatan untuk mendukung proses deteksi dini penyakit diabetes [14]. Berdasarkan kondisi tersebut, masih diperlukan evaluasi terhadap performa Gaussian Naïve Bayes sebagai algoritma klasifikasi yang sederhana pada data kesehatan numerik untuk mengetahui sejauh mana metode tersebut mampu menghasilkan performa yang kompetitif tanpa memerlukan model yang kompleks maupun proses optimasi tambahan.

Berdasarkan *research gap* tersebut, penelitian ini bertujuan untuk mengimplementasikan dan mengevaluasi performa algoritma Gaussian Naïve Bayes dalam mendukung deteksi dini penyakit diabetes menggunakan Pima Indians Diabetes Dataset. Dataset tersebut terdiri atas delapan atribut numerik yang merepresentasikan kondisi kesehatan pasien, yaitu *Pregnancies*, *Glucose*, *Blood Pressure*, *Skin Thickness*, *Insulin*, *Body Mass Index (BMI)*, *Diabetes Pedigree Function*, dan *Age*, sehingga sesuai dengan karakteristik distribusi Gaussian yang digunakan pada algoritma Gaussian Naïve Bayes. Sebelum proses pembentukan model, dilakukan tahap *preprocessing* untuk meningkatkan kualitas data yang digunakan pada proses pelatihan. Selanjutnya, model dibangun menggunakan data pelatihan dan dievaluasi menggunakan data pengujian melalui confusion matrix dengan parameter accuracy, precision, recall, dan F1-score. Penggunaan beberapa parameter evaluasi tersebut bertujuan untuk memberikan gambaran yang lebih komprehensif mengenai kemampuan model dalam mengklasifikasikan pasien ke dalam kategori diabetes maupun non-diabetes, sehingga hasil evaluasi tidak hanya ditinjau berdasarkan tingkat akurasi, tetapi juga kemampuan model dalam mendeteksi kasus positif dan meminimalkan kesalahan klasifikasi.

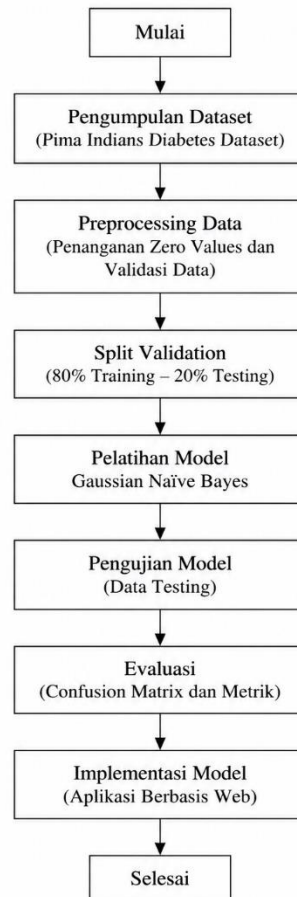
Kontribusi utama penelitian ini adalah memberikan evaluasi terhadap performa Gaussian Naïve Bayes sebagai metode klasifikasi yang sederhana, efisien, dan mudah diimplementasikan untuk mendukung deteksi dini penyakit diabetes. Selain itu, penelitian ini mengintegrasikan model klasifikasi yang telah dibangun ke dalam aplikasi berbasis web sehingga proses prediksi dapat dilakukan secara otomatis berdasarkan data kesehatan yang dimasukkan oleh pengguna. Integrasi tersebut menunjukkan bahwa algoritma dengan kompleksitas komputasi yang relatif rendah masih mampu menghasilkan performa klasifikasi yang baik tanpa memerlukan proses optimasi atau penggunaan model yang lebih kompleks. Dengan demikian, hasil penelitian ini diharapkan dapat menjadi referensi dalam pengembangan sistem pendukung keputusan di bidang kesehatan, khususnya untuk deteksi dini penyakit diabetes, serta menjadi dasar bagi penelitian selanjutnya dalam melakukan evaluasi komparatif menggunakan algoritma klasifikasi lain, teknik validasi yang lebih komprehensif, maupun dataset yang lebih beragam sehingga kualitas model prediksi dapat terus ditingkatkan.

2. METODOLOGI PENELITIAN

2.1 Metodologi Penelitian

Penelitian ini merupakan penelitian kuantitatif dengan pendekatan eksperimen yang bertujuan untuk mengimplementasikan serta mengevaluasi performa algoritma Gaussian Naïve Bayes dalam mendukung deteksi dini penyakit diabetes. Pendekatan eksperimen dipilih karena memungkinkan proses pembentukan, pengujian, dan evaluasi model klasifikasi dilakukan secara terstruktur berdasarkan data kesehatan pasien yang tersedia. Melalui pendekatan ini, model dibangun menggunakan data pelatihan (*training data*) kemudian diuji menggunakan data pengujian (*testing data*) untuk mengetahui kemampuan algoritma dalam mengklasifikasikan pasien ke dalam kategori diabetes maupun non-diabetes. Hasil klasifikasi selanjutnya dianalisis menggunakan beberapa parameter evaluasi sehingga performa model dapat diukur secara objektif dan memberikan gambaran mengenai tingkat ketepatan serta kemampuan model dalam mengenali kasus positif maupun negatif. Selain mengevaluasi performa algoritma, penelitian ini juga

mengimplementasikan model yang telah dibangun ke dalam aplikasi berbasis web sebagai media penerapan hasil klasifikasi sehingga proses prediksi dapat dilakukan secara otomatis berdasarkan data kesehatan yang dimasukkan oleh pengguna. Tahapan penelitian disusun secara sistematis mulai dari proses pengumpulan data, *preprocessing* data, pembentukan model menggunakan algoritma Gaussian Naïve Bayes, pengujian model, evaluasi performa menggunakan confusion matrix, hingga implementasi model ke dalam aplikasi berbasis web [15]. Alur metodologi penelitian yang digunakan pada penelitian ini ditunjukkan pada Gambar 1.



Gambar 1. Metodologi Penelitian

Penelitian menggunakan Pima Indians Diabetes Dataset yang diperoleh dari platform Kaggle, yang merupakan salinan dari UCI Machine Learning Repository. Dataset terdiri atas 768 data pasien dengan delapan atribut masukan, yaitu *Pregnancies*, *Glucose*, *Blood Pressure*, *Skin Thickness*, *Insulin*, *Body Mass Index (BMI)*, *Diabetes Pedigree Function*, dan *Age*, serta satu atribut keluaran (*Outcome*) yang digunakan sebagai kelas klasifikasi. Nilai *Outcome* terdiri atas dua kategori, yaitu 0 yang menunjukkan kondisi non-diabetes dan 1 yang menunjukkan kondisi diabetes. Dataset ini banyak digunakan sebagai acuan dalam penelitian prediksi diabetes karena memiliki karakteristik data numerik yang sesuai untuk penerapan berbagai algoritma klasifikasi [16].

Tahap berikutnya adalah *preprocessing* data yang bertujuan untuk meningkatkan kualitas data sebelum digunakan pada proses pembentukan model. Proses ini meliputi pemeriksaan kelengkapan data, identifikasi *missing values*, pemeriksaan konsistensi atribut, serta penyesuaian tipe data agar sesuai dengan kebutuhan algoritma Gaussian Naïve Bayes. Pada dataset Pima Indians Diabetes, nilai 0 pada atribut *Glucose*, *Blood Pressure*, *Skin Thickness*, *Insulin*, dan *Body Mass Index (BMI)* diperlakukan sebagai *missing values* karena secara fisiologis tidak mungkin bernilai nol. Nilai tersebut kemudian digantikan menggunakan median masing-masing atribut untuk mempertahankan distribusi data dan mengurangi pengaruh *outlier* [17].

Dataset yang telah melalui tahap *preprocessing* selanjutnya dipisahkan menjadi 80% data pelatihan (*training data*) dan 20% data pengujian (*testing data*) menggunakan metode *hold-out validation*. Data pelatihan digunakan untuk membangun model klasifikasi, sedangkan data pengujian digunakan untuk mengevaluasi kemampuan model dalam mengklasifikasikan data yang belum pernah dipelajari sebelumnya. Pemisahan data dilakukan setelah seluruh proses *preprocessing* yang diperlukan agar proses pembentukan model dan evaluasi dapat dilakukan secara sistematis [16].

Model yang telah dibangun kemudian dievaluasi menggunakan confusion matrix dengan parameter accuracy, precision, recall, dan F1-score. Keempat parameter tersebut digunakan untuk memberikan gambaran yang komprehensif mengenai kemampuan model dalam mengklasifikasikan pasien ke dalam kategori diabetes maupun non-diabetes. Selanjutnya, model yang telah diperoleh diintegrasikan ke dalam aplikasi berbasis web sehingga pengguna dapat



memasukkan data kesehatan dan memperoleh hasil prediksi secara otomatis. Pada penelitian ini, implementasi aplikasi berfungsi sebagai media penerapan model klasifikasi, sedangkan fokus utama penelitian tetap diarahkan pada evaluasi performa algoritma Gaussian Naïve Bayes dalam mendukung deteksi dini penyakit diabetes.

2.2 Algoritma Gaussian Naïve Bayes

Algoritma Gaussian Naïve Bayes merupakan salah satu metode klasifikasi probabilistik yang dikembangkan berdasarkan Teorema Bayes dengan asumsi bahwa setiap atribut bersifat saling independen terhadap atribut lainnya. Berbeda dengan Naïve Bayes yang umumnya digunakan pada data kategorikal, Gaussian Naïve Bayes dirancang untuk menangani atribut numerik kontinu dengan memodelkan distribusi setiap atribut menggunakan distribusi normal (Gaussian). Pendekatan ini sesuai diterapkan pada data kesehatan karena sebagian besar atribut, seperti *Glucose*, *Blood Pressure*, *Insulin*, *Body Mass Index (BMI)*, dan *Age*, direpresentasikan dalam bentuk data numerik kontinu [18]. Secara umum, probabilitas suatu data termasuk ke dalam kelas tertentu dihitung menggunakan Teorema Bayes sebagaimana ditunjukkan pada Persamaan (1).

$$P(H | X) = \frac{P(X|H) \times P(H)}{P(X)} \quad (1)$$

Pada Persamaan (1), $P(H | X)$ merupakan probabilitas *posterior*, yaitu peluang suatu data termasuk ke dalam kelas H setelah diketahui atribut X . Nilai $P(X | H)$ merupakan probabilitas *likelihood* yang menunjukkan peluang kemunculan atribut X pada kelas H , sedangkan $P(H)$ merupakan probabilitas awal (*prior probability*) dari kelas H . Adapun $P(X)$ merupakan probabilitas kemunculan atribut (*evidence*) yang berfungsi sebagai faktor normalisasi sehingga total probabilitas seluruh kelas bernilai satu [19].

Karena penelitian ini menggunakan Gaussian Naïve Bayes, maka probabilitas kemunculan setiap atribut numerik dihitung menggunakan fungsi distribusi Gaussian sebagaimana ditunjukkan pada Persamaan (2).

$$P(x_i | C_k) = \frac{1}{\sqrt{2\pi\sigma_k^2}} \exp\left(-\frac{(x_i - \mu_k)^2}{2\sigma_k^2}\right) \quad (2)$$

Pada Persamaan (2), x_i menyatakan nilai atribut ke- i , sedangkan C_k merupakan kelas yang diamati. Simbol μ_k menunjukkan nilai rata-rata (*mean*) atribut pada kelas C_k , sedangkan σ_k merupakan simpangan baku (*standard deviation*) atribut pada kelas tersebut. Nilai $P(x_i | C_k)$ merupakan probabilitas kemunculan atribut x_i terhadap kelas C_k berdasarkan distribusi Gaussian. Probabilitas setiap atribut yang diperoleh dari Persamaan (2) kemudian digunakan sebagai nilai *likelihood* pada Persamaan (1) untuk menghitung probabilitas *posterior* setiap kelas [20].

Proses pembentukan model diawali dengan menghitung probabilitas awal (*prior probability*) setiap kelas berdasarkan data pelatihan. Selanjutnya dihitung nilai *mean* dan *standard deviation* masing-masing atribut pada setiap kelas sebagai parameter distribusi Gaussian. Ketika data baru dimasukkan ke dalam sistem, probabilitas setiap atribut dihitung menggunakan fungsi distribusi Gaussian, kemudian seluruh probabilitas dikombinasikan menggunakan Teorema Bayes untuk memperoleh probabilitas *posterior*. Data selanjutnya diklasifikasikan ke dalam kelas yang memiliki nilai probabilitas *posterior* terbesar sebagai hasil prediksi akhir [21].

Penggunaan Gaussian Naïve Bayes pada penelitian ini dipilih karena memiliki proses komputasi yang sederhana, waktu pelatihan yang relatif cepat, serta tidak memerlukan proses optimasi parameter yang kompleks. Selain itu, algoritma ini mampu memberikan performa klasifikasi yang baik pada data numerik kontinu sehingga sesuai diterapkan sebagai metode deteksi dini penyakit diabetes berdasarkan data kesehatan pasien [21].

2.3 Evaluasi Performa Model

Evaluasi model bertujuan untuk mengukur kemampuan algoritma Gaussian Naïve Bayes dalam mengklasifikasikan data pasien ke dalam kategori diabetes dan non-diabetes. Pada penelitian ini, evaluasi dilakukan menggunakan confusion matrix, yang merupakan salah satu metode evaluasi klasifikasi yang banyak digunakan untuk menganalisis hasil prediksi model berdasarkan jumlah prediksi yang benar maupun salah [22].

Confusion matrix terdiri atas empat komponen utama, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Nilai True Positive menunjukkan jumlah pasien diabetes yang berhasil diprediksi sebagai diabetes, sedangkan True Negative menunjukkan jumlah pasien non-diabetes yang berhasil diprediksi sebagai non-diabetes. Sementara itu, False Positive merupakan jumlah pasien non-diabetes yang diprediksi sebagai diabetes, sedangkan False Negative merupakan jumlah pasien diabetes yang diprediksi sebagai non-diabetes. Keempat komponen tersebut digunakan sebagai dasar dalam menghitung nilai accuracy, precision, recall, dan F1-score [22].

Nilai accuracy digunakan untuk mengukur tingkat ketepatan model dalam mengklasifikasikan seluruh data pengujian. Persamaan accuracy ditunjukkan pada Persamaan (3).

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \times 100\% \quad (3)$$

Pada Persamaan (3), TP (True Positive) merupakan jumlah data positif yang berhasil diklasifikasikan dengan benar, TN (True Negative) merupakan jumlah data negatif yang berhasil diklasifikasikan dengan benar, FP (False Positive) merupakan jumlah data negatif yang salah diklasifikasikan sebagai positif, sedangkan FN (False Negative)



merupakan jumlah data positif yang salah diklasifikasikan sebagai negatif [22]. Nilai precision digunakan untuk mengukur tingkat ketepatan model dalam memberikan prediksi positif. Persamaan precision ditunjukkan pada Persamaan (4).

$$Precision = \frac{TP}{TP+FP} \times 100\% \quad (4)$$

Pada Persamaan (4), TP menunjukkan jumlah prediksi positif yang benar, sedangkan FP menunjukkan jumlah data negatif yang salah diprediksi sebagai positif. Nilai precision yang tinggi menunjukkan bahwa sebagian besar prediksi positif yang dihasilkan model merupakan prediksi yang benar [22]. Nilai recall digunakan untuk mengukur kemampuan model dalam mendeteksi seluruh data positif yang terdapat pada dataset. Persamaan recall ditunjukkan pada Persamaan (5).

$$Recall = \frac{TP}{TP+FN} \times 100\% \quad (5)$$

Pada Persamaan (5), TP merupakan jumlah data positif yang berhasil dikenali dengan benar, sedangkan FN merupakan jumlah data positif yang gagal dikenali oleh model. Nilai recall yang tinggi menunjukkan bahwa model memiliki kemampuan yang baik dalam mendeteksi pasien yang benar-benar menderita diabetes sehingga dapat meminimalkan terjadinya *false negative* [22]. Nilai F1-score digunakan untuk mengukur keseimbangan antara nilai precision dan recall. Persamaan F1-score ditunjukkan pada Persamaan (6).

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \times 100\% \quad (6)$$

Pada Persamaan (6), nilai Precision menunjukkan ketepatan model dalam memberikan prediksi positif, sedangkan Recall menunjukkan kemampuan model dalam mendeteksi seluruh data positif. Nilai F1-score merupakan rata-rata harmonik dari precision dan recall sehingga mampu memberikan gambaran yang lebih seimbang terhadap performa model, terutama pada permasalahan klasifikasi yang memiliki distribusi kelas tidak sepenuhnya seimbang [22].

Keempat parameter evaluasi tersebut digunakan secara bersama-sama untuk memberikan gambaran yang komprehensif mengenai performa algoritma Gaussian Naïve Bayes. Penggunaan beberapa metrik evaluasi memungkinkan analisis yang lebih objektif dibandingkan hanya menggunakan nilai accuracy, karena setiap metrik memberikan informasi yang berbeda mengenai kemampuan model dalam melakukan klasifikasi data pasien.

3. HASIL DAN PEMBAHASAN

3.1 Hasil Pengolahan Dataset

Tahap awal penelitian dilakukan melalui pengolahan *Pima Indians Diabetes Dataset* yang diperoleh dari Kaggle. Dataset terdiri atas 768 data pasien dengan delapan atribut masukan, yaitu *Pregnancies*, *Glucose*, *Blood Pressure*, *Skin Thickness*, *Insulin*, *Body Mass Index (BMI)*, *Diabetes Pedigree Function*, dan *Age*, serta satu atribut keluaran (*Outcome*) yang digunakan untuk mengklasifikasikan pasien ke dalam kategori diabetes dan non-diabetes. Sebelum digunakan dalam pembentukan model Naïve Bayes, dataset terlebih dahulu melalui tahap *preprocessing* yang meliputi pemeriksaan kelengkapan, konsistensi, dan penyesuaian format data guna meningkatkan kualitas data serta meminimalkan kesalahan klasifikasi. Proses ini dilakukan untuk memastikan bahwa data yang digunakan layak dan siap diproses oleh algoritma sehingga dapat menghasilkan model yang lebih optimal. Selain itu, tahap *preprocessing* juga membantu mengurangi potensi kesalahan yang dapat memengaruhi hasil prediksi. Setelah tahap *preprocessing* selesai, data kemudian dibagi menjadi data pelatihan (*training data*) dan data pengujian (*testing data*) menggunakan metode *split validation* untuk mendukung proses pembentukan dan evaluasi model klasifikasi. Hasil pengolahan dataset tersebut menjadi dasar dalam proses pelatihan dan pengujian algoritma Naïve Bayes pada penelitian ini.

3.1.1 Dataset Penelitian

Dataset yang digunakan dalam penelitian ini terdiri atas 768 data pasien dengan sembilan atribut yang digunakan dalam proses klasifikasi. Ringkasan dataset penelitian ditunjukkan pada Tabel 4.

Tabel 4. Ringkasan Dataset Penelitian

No	Keterangan	Jumlah
1	Total Data	768
2	Atribut Input	8
3	Atribut Output	1
4	Jumlah Kelas	2
5	Kategori Output	Diabetes dan Non-Diabetes

Berdasarkan Tabel 4, dataset yang digunakan memiliki dua kategori kelas yang akan diprediksi oleh sistem, yaitu kelas diabetes dan non-diabetes. Seluruh atribut masukan digunakan sebagai parameter dalam proses pembentukan model Naïve Bayes untuk menentukan kemungkinan seorang pasien termasuk ke dalam salah satu kelas tersebut. Dataset ini

dipilih karena memiliki atribut yang relevan dengan faktor risiko diabetes dan telah banyak digunakan sebagai acuan dalam penelitian klasifikasi penyakit diabetes. Jumlah data yang tersedia juga dinilai cukup untuk mendukung proses pelatihan dan pengujian model sehingga hasil klasifikasi yang diperoleh dapat digunakan sebagai dasar dalam mengevaluasi performa algoritma yang diterapkan.

3.1.2 Hasil Pembagian Dataset

Setelah proses *preprocessing* selesai dilakukan, dataset dibagi menjadi data pelatihan (*training data*) dan data pengujian (*testing data*) menggunakan metode *split validation*. Pada penelitian ini digunakan rasio pembagian data sebesar 80:20, di mana 80% data digunakan sebagai data pelatihan dan 20% data digunakan sebagai data pengujian.

Pembagian dataset dilakukan untuk memastikan bahwa model dapat mempelajari pola data dari data pelatihan dan selanjutnya diuji menggunakan data yang belum pernah dipelajari sebelumnya. Dengan metode ini, kemampuan model dalam melakukan generalisasi terhadap data baru dapat dievaluasi secara lebih objektif.

Hasil pembagian dataset ditunjukkan pada Tabel 5.

Tabel 5. Hasil Pembagian Dataset

No	Jenis Data	Jumlah Data	Persentase
1	Training Data	614	80%
2	Testing Data	154	20%
	Total	768	100%

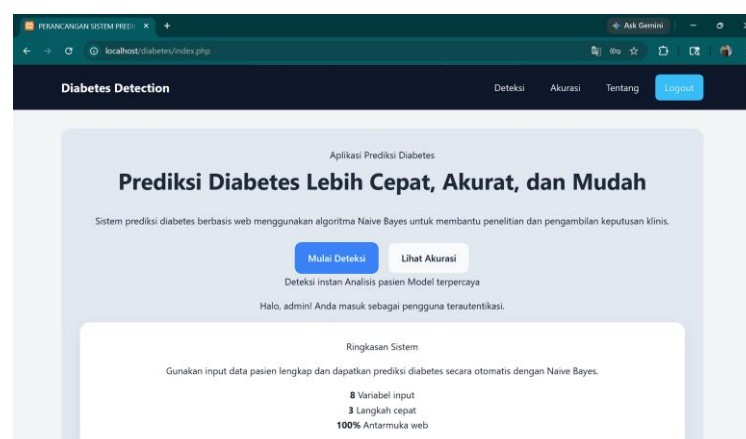
Berdasarkan Tabel 5, sebanyak 614 data digunakan sebagai data pelatihan untuk membangun model klasifikasi Naïve Bayes, sedangkan 154 data digunakan sebagai data pengujian untuk mengevaluasi performa model. Jumlah data pelatihan yang lebih besar dibandingkan data pengujian bertujuan agar model memiliki cukup informasi untuk mempelajari pola hubungan antaratribut yang terdapat pada dataset. Hasil pengolahan dataset menunjukkan bahwa seluruh data berhasil dipersiapkan dan dapat digunakan dalam proses pembentukan model klasifikasi. Dataset yang telah melalui tahap *preprocessing* dan pembagian data selanjutnya digunakan pada proses pelatihan model Naïve Bayes serta pengujian performa model yang akan dibahas pada subbab berikutnya.

3.2 Implementasi Sistem Berbasis Web

Implementasi sistem merupakan tahap penerapan model klasifikasi Naïve Bayes ke dalam aplikasi berbasis web yang dikembangkan untuk membantu proses prediksi penyakit diabetes. Sistem dibangun menggunakan bahasa pemrograman PHP dan basis data MySQL sehingga dapat diakses melalui peramban web. Implementasi sistem meliputi beberapa halaman utama yang digunakan oleh pengguna untuk melakukan deteksi diabetes, melihat hasil pengujian model, serta memperoleh informasi mengenai sistem yang dikembangkan. Pengembangan sistem berbasis web dilakukan untuk memberikan kemudahan akses kepada pengguna tanpa memerlukan instalasi perangkat lunak tambahan pada perangkat yang digunakan. Selain itu, sistem dirancang dengan antarmuka yang sederhana dan mudah dipahami sehingga dapat digunakan oleh berbagai kalangan pengguna. Melalui implementasi ini, proses prediksi risiko diabetes dapat dilakukan secara lebih cepat dan efisien dengan memanfaatkan hasil klasifikasi yang dihasilkan oleh algoritma Naïve Bayes. Dengan adanya sistem yang terkomputerisasi, pengguna dapat memperoleh informasi awal mengenai kemungkinan risiko diabetes berdasarkan data kesehatan yang dimasukkan ke dalam sistem.

3.2.1 Implementasi Halaman Beranda

Halaman beranda merupakan halaman awal yang ditampilkan ketika pengguna mengakses sistem. Halaman ini berfungsi sebagai media informasi utama yang menjelaskan tujuan sistem serta menyediakan navigasi menuju fitur-fitur yang tersedia. Selain itu, halaman beranda juga memberikan gambaran singkat mengenai sistem prediksi diabetes yang dibangun menggunakan algoritma Naïve Bayes.

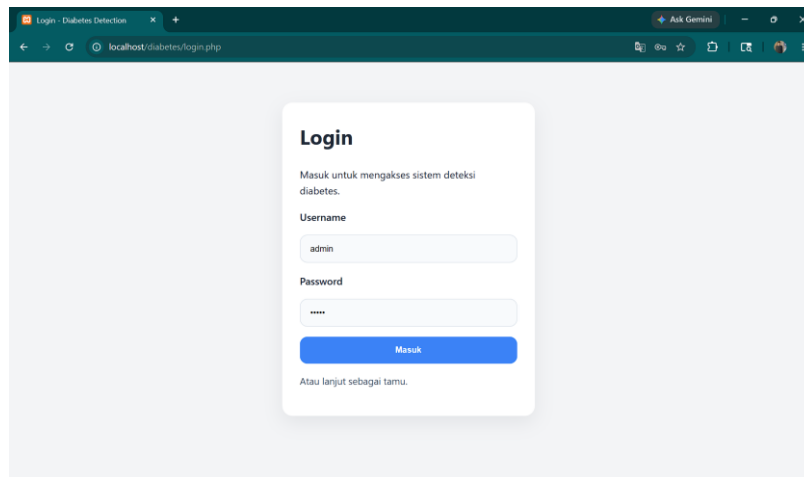


Gambar 3. Halaman Beranda

Berdasarkan Gambar 3, halaman beranda dirancang dengan tampilan yang sederhana dan mudah dipahami oleh pengguna. Melalui halaman ini, pengguna dapat memperoleh informasi awal mengenai sistem sebelum melakukan proses deteksi diabetes.

3.2.2 Implementasi Halaman Login

Halaman login digunakan sebagai proses autentikasi pengguna sebelum mengakses fitur utama sistem. Pada halaman ini pengguna diminta memasukkan *username* dan *password* yang telah terdaftar pada sistem. Proses login bertujuan untuk menjaga keamanan data dan membatasi akses hanya kepada pengguna yang memiliki hak akses.

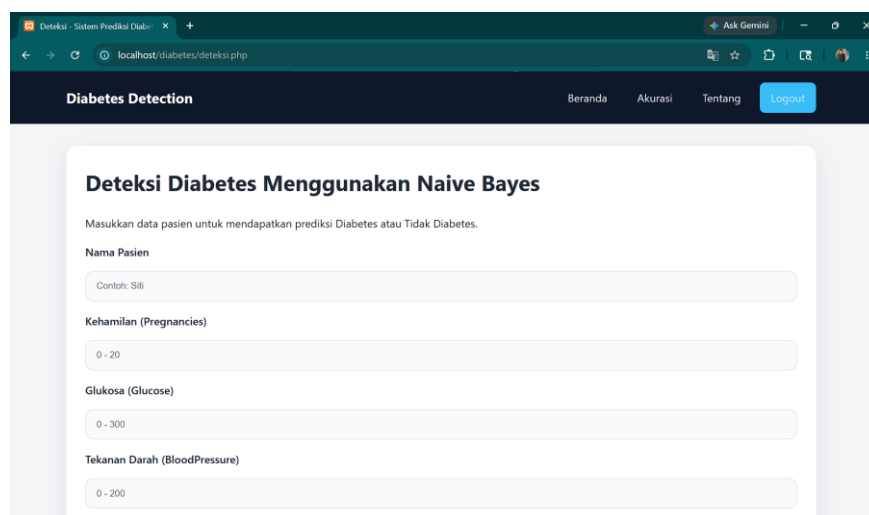


Gambar 4. Halaman Login

Berdasarkan Gambar 4, halaman login menyediakan form autentikasi yang terdiri atas kolom *username* dan *password*. Setelah data yang dimasukkan sesuai, pengguna akan diarahkan menuju halaman utama sistem.

3.2.3 Implementasi Halaman Deteksi Diabetes

Halaman deteksi diabetes merupakan fitur utama yang digunakan untuk melakukan proses prediksi. Pada halaman ini pengguna dapat memasukkan data kesehatan yang diperlukan sebagai parameter klasifikasi, seperti jumlah kehamilan (*Pregnancies*), kadar glukosa (*Glucose*), tekanan darah (*Blood Pressure*), ketebalan kulit (*Skin Thickness*), kadar insulin (*Insulin*), indeks massa tubuh (*Body Mass Index*), *Diabetes Pedigree Function*, dan usia (*Age*).

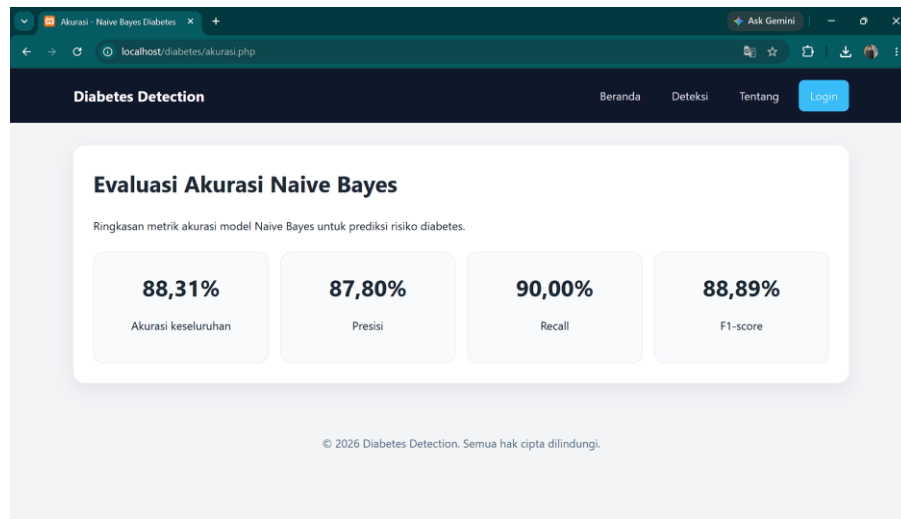


Gambar 5. Halaman Deteksi Diabetes

Berdasarkan Gambar 5, sistem menyediakan formulir input yang digunakan untuk memasukkan data kesehatan pasien. Setelah seluruh data diisi, sistem akan melakukan proses klasifikasi menggunakan algoritma Naïve Bayes dan menghasilkan prediksi berupa kategori diabetes atau non-diabetes.

3.2.4 Implementasi Halaman Akurasi

Halaman akurasi digunakan untuk menampilkan hasil pengujian model yang telah dilakukan menggunakan data pengujian (*testing data*). Informasi yang ditampilkan meliputi nilai *accuracy*, *precision*, *recall*, dan *F1-score* yang diperoleh dari evaluasi menggunakan *confusion matrix*.

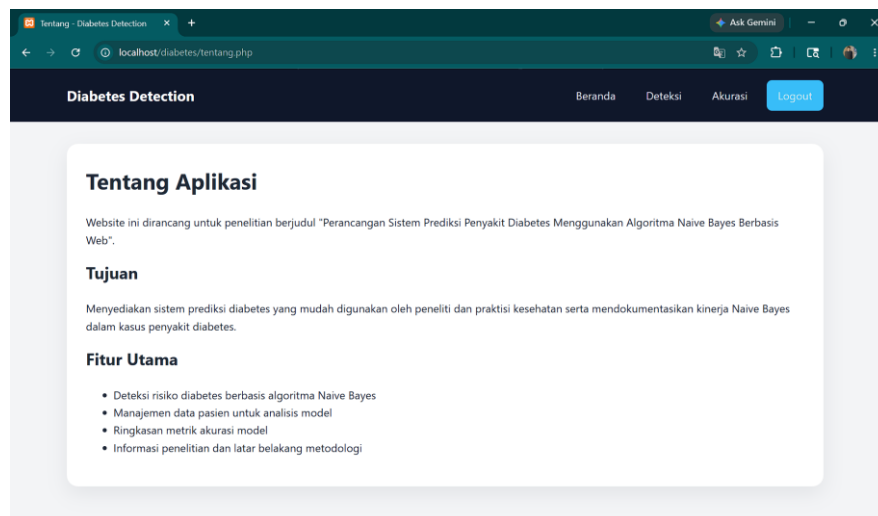


Gambar 6. Halaman Akurasi

Berdasarkan Gambar 6, halaman akurasi memberikan informasi mengenai performa model Naïve Bayes yang digunakan dalam sistem. Informasi tersebut digunakan untuk mengetahui tingkat ketepatan model dalam melakukan klasifikasi terhadap data pasien.

3.2.5 Implementasi Halaman Tentang

Halaman tentang (*about*) berisi informasi mengenai sistem yang dikembangkan, tujuan penelitian, serta teknologi yang digunakan dalam pengembangan aplikasi. Halaman ini bertujuan untuk memberikan informasi tambahan kepada pengguna mengenai sistem prediksi diabetes berbasis web.



Gambar 7. Halaman Tentang

Berdasarkan Gambar 7, halaman tentang memberikan informasi mengenai latar belakang pengembangan sistem serta metode yang digunakan dalam proses prediksi diabetes. Informasi tersebut membantu pengguna memahami fungsi dan tujuan sistem yang digunakan.

Implementasi sistem berbasis web yang telah dikembangkan menunjukkan bahwa seluruh fitur dapat berjalan sesuai dengan fungsinya masing-masing. Sistem mampu menerima data masukan dari pengguna, melakukan proses klasifikasi menggunakan algoritma Naïve Bayes, serta menampilkan hasil prediksi dan informasi performa model secara otomatis. Dengan demikian, sistem yang dihasilkan dapat dimanfaatkan sebagai alat bantu dalam proses deteksi dini risiko penyakit diabetes secara cepat, mudah, dan efisien.

3.3 Hasil Pengujian Model Naïve Bayes

Pengujian model dilakukan untuk mengevaluasi performa algoritma Gaussian Naïve Bayes dalam mengklasifikasikan data pasien ke dalam kategori diabetes dan non-diabetes. Evaluasi dilakukan menggunakan 154 data pengujian yang diperoleh melalui metode *split validation* dengan rasio 80% data pelatihan dan 20% data pengujian. Seluruh data pengujian merupakan data yang tidak digunakan selama proses pelatihan sehingga hasil evaluasi dapat menggambarkan kemampuan model dalam melakukan klasifikasi terhadap data baru.



Performa model diukur menggunakan metode confusion matrix, yang kemudian digunakan sebagai dasar dalam menghitung nilai accuracy, precision, recall, dan F1-score. Penggunaan beberapa parameter evaluasi tersebut bertujuan untuk memberikan gambaran yang lebih komprehensif mengenai kemampuan model dalam melakukan klasifikasi dibandingkan hanya menggunakan nilai *accuracy*.

3.3.1 Hasil Confusion Matrix

Hasil klasifikasi menggunakan algoritma Gaussian Naïve Bayes dievaluasi menggunakan confusion matrix untuk mengetahui jumlah prediksi yang benar maupun prediksi yang salah pada setiap kelas. Hasil confusion matrix ditunjukkan pada Tabel 6.

Tabel 6. Hasil Confusion Matrix

Actual Class	Prediksi Diabetes	Prediksi Non-Diabetes
Diabetes	72 (TP)	8 (FN)
Non-Diabetes	10 (FP)	64 (TN)

Berdasarkan Tabel 6, diperoleh True Positive (TP) sebanyak 72 data, yang menunjukkan jumlah pasien diabetes yang berhasil diprediksi sebagai diabetes. Nilai True Negative (TN) sebanyak 64 data menunjukkan jumlah pasien non-diabetes yang berhasil diprediksi sebagai non-diabetes. Selain itu, terdapat False Positive (FP) sebanyak 10 data, yaitu pasien non-diabetes yang diprediksi sebagai diabetes, sedangkan False Negative (FN) sebanyak 8 data, yaitu pasien diabetes yang diprediksi sebagai non-diabetes.

Secara keseluruhan, dari 154 data pengujian, model berhasil mengklasifikasikan 136 data secara benar dan 18 data mengalami kesalahan klasifikasi. Hasil tersebut menunjukkan bahwa Gaussian Naïve Bayes mampu mempelajari pola pada data kesehatan pasien dengan baik sehingga menghasilkan tingkat klasifikasi yang tinggi. Selain itu, jumlah False Negative yang lebih sedikit dibandingkan False Positive mengindikasikan bahwa model memiliki sensitivitas yang baik dalam mendeteksi pasien yang benar-benar menderita diabetes, sehingga sesuai untuk mendukung proses deteksi dini penyakit diabetes.

3.3.2 Hasil Evaluasi Model

Berdasarkan confusion matrix, dilakukan perhitungan terhadap parameter accuracy, precision, recall, dan F1-score. Rekapitulasi hasil evaluasi model ditunjukkan pada Tabel 7.

Tabel 7. Rekapitulasi Hasil Evaluasi Model

Parameter	Nilai
Accuracy	88,31%
Precision	87,80%
Recall	90,00%
F1-Score	88,89%

Berdasarkan Tabel 7 diperoleh nilai accuracy sebesar 88,31%, yang menunjukkan bahwa model mampu mengklasifikasikan sebagian besar data pengujian secara benar. Nilai tersebut mengindikasikan bahwa algoritma Gaussian Naïve Bayes memiliki tingkat ketepatan klasifikasi yang baik dalam membedakan pasien diabetes dan non-diabetes berdasarkan atribut kesehatan yang digunakan.

Nilai precision sebesar 87,80% menunjukkan bahwa sebagian besar pasien yang diprediksi sebagai penderita diabetes memang benar-benar termasuk ke dalam kategori diabetes. Sementara itu, nilai recall sebesar 90,00% menunjukkan bahwa model memiliki kemampuan yang lebih tinggi dalam mendeteksi pasien yang benar-benar menderita diabetes dibandingkan mempertahankan ketepatan seluruh prediksi positif.

Perbedaan nilai recall yang sedikit lebih tinggi dibandingkan precision menunjukkan bahwa model menghasilkan jumlah False Negative yang lebih sedikit dibandingkan False Positive. Pada penelitian ini diperoleh False Negative sebanyak 8 data, sedangkan False Positive sebanyak 10 data. Kondisi tersebut menunjukkan bahwa model lebih sensitif dalam mendeteksi pasien yang berisiko menderita diabetes sehingga kemungkinan pasien diabetes yang tidak terdeteksi menjadi lebih kecil. Dalam konteks deteksi dini penyakit diabetes, karakteristik tersebut merupakan kondisi yang diharapkan karena kesalahan berupa False Negative memiliki konsekuensi yang lebih besar dibandingkan False Positive. Pasien yang tidak terdeteksi sebagai penderita diabetes berpotensi mengalami keterlambatan diagnosis dan penanganan sehingga meningkatkan risiko terjadinya komplikasi. Sebaliknya, pasien yang mengalami False Positive masih dapat menjalani pemeriksaan medis lanjutan untuk memastikan kondisi kesehatannya.

Nilai F1-score sebesar 88,89% menunjukkan bahwa model memiliki keseimbangan yang baik antara precision dan recall. Hal tersebut mengindikasikan bahwa algoritma Gaussian Naïve Bayes tidak hanya mampu memberikan prediksi yang tepat, tetapi juga memiliki sensitivitas yang baik dalam mengidentifikasi pasien yang berisiko menderita diabetes.

Secara keseluruhan, hasil evaluasi menunjukkan bahwa algoritma Gaussian Naïve Bayes memiliki performa yang baik dalam melakukan klasifikasi penyakit diabetes. Kombinasi nilai accuracy, precision, recall, dan F1-score yang seluruhnya berada di atas 87% menunjukkan bahwa model mampu memberikan hasil klasifikasi yang konsisten dan dapat digunakan sebagai dasar dalam mendukung proses deteksi dini penyakit diabetes.



3.4 Pembahasan

Hasil penelitian menunjukkan bahwa algoritma Gaussian Naïve Bayes mampu memberikan performa klasifikasi yang baik dalam mendukung deteksi dini penyakit diabetes. Berdasarkan hasil pengujian diperoleh nilai accuracy sebesar 88,31%, precision sebesar 87,80%, recall sebesar 90,00%, dan F1-score sebesar 88,89%. Nilai tersebut menunjukkan bahwa model memiliki tingkat ketepatan klasifikasi yang tinggi serta mampu mempertahankan keseimbangan antara ketepatan prediksi dan kemampuan mendeteksi pasien yang berisiko diabetes.

Salah satu temuan penting dalam penelitian ini adalah nilai recall yang lebih tinggi dibandingkan precision. Kondisi tersebut menunjukkan bahwa model lebih sensitif dalam mengidentifikasi pasien yang benar-benar menderita diabetes dibandingkan menjaga ketepatan seluruh prediksi positif. Berdasarkan hasil *confusion matrix*, jumlah False Negative lebih sedikit dibandingkan False Positive, sehingga kemungkinan pasien diabetes yang tidak terdeteksi menjadi lebih rendah. Dalam konteks deteksi dini, karakteristik tersebut lebih menguntungkan karena kesalahan berupa False Negative dapat menyebabkan keterlambatan diagnosis dan penanganan medis yang berpotensi meningkatkan risiko komplikasi penyakit. Sebaliknya, pasien yang termasuk dalam kategori False Positive masih dapat menjalani pemeriksaan lanjutan untuk memastikan diagnosis sehingga risiko klinis yang ditimbulkan relatif lebih kecil.

Performa yang diperoleh menunjukkan bahwa Gaussian Naïve Bayes tetap mampu menghasilkan klasifikasi yang baik meskipun menggunakan pendekatan probabilistik yang sederhana tanpa penerapan teknik optimasi maupun seleksi fitur. Penggunaan distribusi Gaussian sesuai dengan karakteristik atribut numerik pada dataset sehingga memungkinkan model menghitung probabilitas setiap atribut secara efektif. Selain memiliki proses pelatihan yang cepat, algoritma ini juga memiliki kompleksitas komputasi yang relatif rendah sehingga sesuai diterapkan pada aplikasi pendukung deteksi dini yang memerlukan proses prediksi secara efisien. Untuk memberikan gambaran posisi hasil penelitian terhadap penelitian sebelumnya, dilakukan perbandingan dengan beberapa penelitian yang menggunakan algoritma klasifikasi pada kasus prediksi diabetes sebagaimana ditunjukkan pada Tabel 8.

Tabel 8. Perbandingan Penelitian Terdahulu dengan Penelitian Ini

Penelitian	Metode	Temuan Utama
Resti et al. [6]	Naïve Bayes	Menunjukkan bahwa Naïve Bayes efektif untuk diagnosis diabetes mellitus.
Salsabila et al. [7]	Naïve Bayes	Naïve Bayes memberikan performa klasifikasi yang baik pada prediksi diabetes.
Amal et al. [8]	Naïve Bayes	Pendekatan probabilistik mampu mengidentifikasi risiko diabetes secara efektif.
Islam et al. [9]	Stacked Machine Learning & Deep Learning	Algoritma yang lebih kompleks memberikan performa tinggi, namun membutuhkan komputasi yang lebih besar.
Penelitian ini	Gaussian Naïve Bayes	Menghasilkan accuracy 88,31%, precision 87,80%, recall 90,00%, dan F1-score 88,89% dengan implementasi model yang sederhana dan efisien.

Berdasarkan Tabel 8, hasil penelitian ini menunjukkan bahwa algoritma Gaussian Naïve Bayes memperoleh accuracy sebesar 88,31%, sehingga sejalan dengan penelitian Resti et al. yang menunjukkan bahwa Naïve Bayes efektif digunakan dalam diagnosis diabetes mellitus [6]. Temuan ini juga konsisten dengan penelitian Salsabila et al. yang melaporkan bahwa Naïve Bayes mampu memberikan performa klasifikasi yang baik pada prediksi diabetes [7]. Selain itu, Amal et al. menyimpulkan bahwa pendekatan probabilistik masih efektif dalam mengidentifikasi risiko diabetes berdasarkan data kesehatan pasien [8]. Meskipun Islam et al. melaporkan bahwa penggunaan *stacked machine learning* dan *deep learning* mampu menghasilkan performa klasifikasi yang lebih tinggi [9], hasil penelitian ini menunjukkan bahwa Gaussian Naïve Bayes tetap mampu memberikan performa yang kompetitif dengan kompleksitas komputasi yang lebih rendah serta proses implementasi yang lebih sederhana.

Performa tersebut menunjukkan bahwa penggunaan algoritma yang lebih kompleks tidak selalu menjadi pilihan yang paling tepat, terutama pada kasus klasifikasi dengan dataset berukuran relatif kecil seperti Pima Indians Diabetes Dataset. Dengan karakteristik data yang didominasi oleh atribut numerik kontinu, Gaussian Naïve Bayes mampu memodelkan distribusi data secara efektif sehingga menghasilkan tingkat akurasi, *precision*, *recall*, dan *F1-score* yang baik tanpa memerlukan proses optimasi tambahan. Dari sisi implementasi, pendekatan ini juga lebih mudah diintegrasikan ke dalam aplikasi berbasis web karena memiliki proses pelatihan dan prediksi yang cepat serta kebutuhan komputasi yang relatif rendah. Oleh karena itu, Gaussian Naïve Bayes dapat dipertimbangkan sebagai alternatif metode klasifikasi yang efisien untuk mendukung proses deteksi dini penyakit diabetes.

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa algoritma Gaussian Naïve Bayes mampu memberikan performa klasifikasi yang baik dalam mendukung deteksi dini penyakit diabetes berdasarkan data kesehatan pasien. Hasil evaluasi menunjukkan bahwa model memperoleh accuracy sebesar 88,31%, precision sebesar 87,80%, recall sebesar 90,00%, dan F1-score sebesar 88,89%, yang menunjukkan keseimbangan antara ketepatan prediksi dan



kemampuan mendeteksi pasien yang berisiko diabetes. Nilai recall yang lebih tinggi dibandingkan precision menunjukkan bahwa model lebih sensitif dalam mengidentifikasi kasus positif sehingga berpotensi mengurangi kemungkinan pasien diabetes tidak terdeteksi pada tahap skrining awal. Temuan ini menunjukkan bahwa Gaussian Naïve Bayes tidak hanya memiliki performa klasifikasi yang baik, tetapi juga menawarkan proses komputasi yang sederhana dan efisien sehingga layak diterapkan sebagai model pendukung pada aplikasi deteksi dini penyakit diabetes. Meskipun penelitian ini hanya menggunakan satu dataset publik dan belum melakukan perbandingan eksperimental dengan algoritma klasifikasi lain, hasil yang diperoleh menunjukkan bahwa pendekatan probabilistik sederhana masih relevan untuk diterapkan pada kasus klasifikasi di bidang kesehatan. Oleh karena itu, penelitian selanjutnya disarankan menggunakan dataset yang lebih beragam, menerapkan metode validasi yang lebih komprehensif, serta membandingkan performa Gaussian Naïve Bayes dengan algoritma klasifikasi lainnya untuk memperoleh model yang lebih optimal.

REFERENCES

- [1] D. J. Magliano and E. J. Boyko, *IDF Diabetes Atlas*, 11th ed. Brussels, Belgium: International Diabetes Federation, 2025. [Online]. Available: <https://diabetesatlas.org>
- [2] N. A. ElSayed et al., "Diagnosis and Classification of Diabetes: Standards of Care in Diabetes—2025," *Diabetes Care*, vol. 48, no. Supplement_1, pp. S27–S49, 2025, doi: 10.2337/dc25-S002.
- [3] M. Kiran, Y. Xie, N. Anjum, G. Ball, B. Pierscionek, and D. Russell, "Machine Learning and Artificial Intelligence in Type 2 Diabetes Prediction: A Comprehensive 33-Year Bibliometric and Literature Analysis," *Frontiers in Digital Health*, vol. 7, 2025, doi: 10.3389/fdgth.2025.1557467.
- [4] P. D. Petridis, A. S. Kristo, A. K. Sikilidis, and I. K. Kitsas, "A Review on Trending Machine Learning Techniques for Type 2 Diabetes Mellitus Management," *Informatics*, vol. 11, no. 4, Art. no. 70, 2024, doi: 10.3390/informatics11040070.
- [5] S. Afolabi, N. Ajadi, A. Jimoh, and I. Adenekan, "Predicting Diabetes Using Supervised Machine Learning Algorithms on E-Health Records," *Informatics and Health*, vol. 2, no. 1, pp. 9–16, 2025, doi: 10.1016/j.infoh.2024.12.002.
- [6] Y. Resti, E. S. Kresnawati, N. R. Dewi, D. A. Zayanti, and N. Eliyati, "Diagnosis of Diabetes Mellitus in Women of Reproductive Age Using the Prediction Methods of Naive Bayes, Discriminant Analysis, and Logistic Regression," *Science and Technology Indonesia*, vol. 6, no. 2, pp. 96–104, 2021, doi: 10.26554/sti.2021.6.2.96-104.
- [7] L. N. Salsabila, M. R. Dwi Pangga, S. M. Yasser, N. A. Riyani, S. Aminah, and W. Wahyunengsih, "Application of Naïve Bayes Algorithm for Diabetes Prediction," *Unisda Journal of Mathematics and Computer Science (UJMC)*, vol. 10, no. 1, pp. 56–68, 2024, doi: 10.52166/ujmc.v10i1.6886.
- [8] L. H. I. Amal et al., "Performance Analysis of Naive Bayes Method for Diabetes Diagnosis," *International Journal of Healthcare and Information Technology*, vol. 3, no. 2, pp. 78–87, 2025, doi: 10.25047/ijhitech.v3i2.6670.
- [9] M. Z. Islam, "Enhancing Diabetes Prediction Accuracy Using Stacked Machine Learning and Deep Learning Models: A Public Health Approach," *The Indonesian Journal of Computer Science*, vol. 14, no. 4, 2025, doi: 10.33022/ijcs.v14i4.4947.
- [10] H. V. Nguyen, Y. Choi, and H. Byeon, "An Explainable Hybrid Deep Learning Model for Prediabetes Prediction in Men Aged 30 and Above," *Journal of Men's Health*, vol. 20, no. 10, Art. no. 166, 2024, doi: 10.22514/jomh.2024.166.
- [11] S. Majyambere, T. Lindgren, C. Twizere, and I. Ntakirutimana, "Early Type 2 Diabetes Risk Prediction Using Explainable Machine Learning in a Two-Stage Approach," *Frontiers in Digital Health*, vol. 8, 2026, doi: 10.3389/fdgth.2026.1743619.
- [12] M. S. Alzaboon, M. Alqaraleh, and M. S. Al-Batah, "Diabetes Prediction and Management Using Machine Learning Approaches," *Data and Metadata*, vol. 4, Art. no. 545, 2025, doi: 10.56294/dm2025545.
- [13] D. S. Khafaga, A. H. Alharbi, I. Mohamed, and K. M. Hosny, "An Integrated Classification and Association Rule Technique for Early-Stage Diabetes Risk Prediction," *Healthcare*, vol. 10, no. 10, Art. no. 2070, 2022, doi: 10.3390/healthcare10102070.
- [14] A. H. Mousa et al., "Diabetes at a Glance: Assessing AI Strategies for Early Diabetes Detection and Intervention via a Mobile App," *Mesopotamian Journal of Computer Science*, vol. 2025, pp. 288–301, 2025, doi: 10.58496/MJCSC/2025/018.
- [15] C. C. Olisah, L. Smith, and M. Smith, "Diabetes Mellitus Prediction and Diagnosis from a Data Preprocessing and Machine Learning Perspective," *Computer Methods and Programs in Biomedicine*, vol. 220, Art. no. 106773, 2022, doi: 10.1016/j.cmpb.2022.106773.
- [16] A. M. AbdulAbbas, R. Alkanany, Y. A. K. Al-Nuaimi, and Z. M. A. Al-Hamdawee, "A Sequential Data Preprocessing Pipeline for Diabetes Prediction: A Data Leakage Prevention and Dual-Validation Approach," *Engineering, Technology & Applied Science Research*, vol. 15, no. 6, pp. 30059–30066, 2025, doi: 10.48084/etasr.14155.
- [17] V. Jain, S. Shukla, and N. Khare, "Analysis of Various Data Imputation Techniques for Diabetes Classification on PIMA Dataset," in *2024 IEEE International Students' Conference on Electrical, Electronics and Computer Science (SCEECs)*, Piscataway, NJ, USA: IEEE, 2024, pp. 1–6, doi: 10.1109/SCEECs61402.2024.10482050.
- [18] V. Jaiswal, A. Negi, and T. Pal, "A Review on Current Advances in Machine Learning Based Diabetes Prediction," *Primary Care Diabetes*, vol. 15, no. 3, pp. 435–443, 2021, doi: 10.1016/j.pcd.2021.02.005.
- [19] K. B. Lesmana and I. K. G. Suhartana, "Deteksi Penyakit Diabetes Menggunakan Gaussian Naive Bayes, Regresi Logistik, dan Random Forest," *Jurnal Nasional Teknologi Informasi dan Aplikasinya*, vol. 1, no. 4, pp. 1209–1214, 2023, doi: 10.24843/JNATIA.2023.v01.i04.p25.
- [20] Y. Mao et al., "Value of Machine Learning Algorithms for Predicting Diabetes Risk: A Subset Analysis from a Real-World Retrospective Cohort Study," *Journal of Diabetes Investigation*, vol. 14, no. 2, pp. 309–320, 2023, doi: 10.1111/jdi.13937.
- [21] S. A. Hicks et al., "On Evaluation Metrics for Medical Applications of Artificial Intelligence," *Scientific Reports*, vol. 12, no. 1, Art. no. 5979, 2022, doi: 10.1038/s41598-022-09954-8.