

PERBANDINGAN ALGORITMA KNN DAN NAIVE BAYES PADA DATA BALITA UNTUK KLASIFIKASI STATUS GIZI

Fahmi Djuniar, Risqiati, Arief Soma Darmawan

Teknik Informatika, Institut Widya Pratama

Jl. Patriot No. 25 - Kota Pekalongan

fahmidjun@gmail.com

ABSTRAK

Permasalahan gizi balita, khususnya stunting, merupakan ancaman serius bagi generasi penerus bangsa. Berdasarkan Survei Kesehatan Indonesia (SKI) tahun 2023, prevalensi stunting di Indonesia masih berada di angka 21,5% — jauh di atas ambang kritis WHO sebesar 20% dan belum mencapai target 14% yang diamanatkan Peraturan Presiden Nomor 72 Tahun 2021 tentang Percepatan Penurunan Stunting (Kementerian Kesehatan RI, 2024). Kondisi ini menuntut solusi deteksi dini yang cepat, akurat, dan berbasis data. Pemilihan algoritma machine learning yang tepat sangat menentukan efektivitas sistem klasifikasi status gizi tersebut. Meskipun K-Nearest Neighbor (KNN) dan Naive Bayes Classifier (NBC) telah banyak digunakan dalam penelitian serupa, hasil yang diperoleh bervariasi tergantung pada karakteristik data lokal sehingga perbandingan langsung keduanya pada dataset Puskesmas Tirto, Kota Pekalongan, menjadi penting untuk dilakukan. Penelitian ini dilakukan untuk membandingkan 2 algoritma yang berbeda yaitu K-Nearest Neighbors (KNN) dan Naive Bayes Classifier (NBC), data penelitian ini di dapat dari Puskesmas Tirto Kota Pekalongan. Status gizi ditentukan berdasarkan indikator pengukuran tubuh manusia yang meliputi BB, TB, JK, dan usia. Penelitian ini memakai metode data hold out dengan membagi datanya 80% sebagai data latih dan 20% sebagai data testing. Pengukuran performa model dilakukan menggunakan indikator matrik Akurasi, Presisi, Recall, dan F1-score yang didapat dari confusion matrik. Hasil pengujian mendapati bahwa algoritma KNN mendapatkan performa lebih bagus dari pada NBC dengan nilai Accuracy 85,37%, Precision 80%, Recall 85%, dan F1-score 83%. sedangkan itu, algoritma Naive Bayes Classifier (NBC) memperoleh nilai Accuracy 68,29%, dengan Precision 79%, Recall 68%, dan F1-score 73%, ini menunjukkan bahwa performa lebih rendah terutama pada klasifikasi kelas minoritas. jadi algoritma KNN menjadi yang paling optimal dalam melakukan klasifikasi status gizi balita pada dataset penelitian ini.

Kata kunci : *Classification, K-Nearest Neighbors (KNN), Naive Bayes Classifier (NBC), Nutritional Status*

1. PENDAHULUAN

Masalah gizi pada balita, khususnya stunting (tengkes), masih menjadi tantangan serius dalam bidang kesehatan masyarakat di Indonesia. Stunting merupakan kondisi gagal tumbuh pada anak balita akibat kekurangan gizi kronis yang berdampak pada perkembangan fisik dan kognitif jangka panjang. Berdasarkan Survei Status Gizi Indonesia (SSGI) 2022 yang diselenggarakan Kementerian Kesehatan RI, prevalensi stunting nasional pada balita tercatat sebesar 21,6%, turun dari 24,4% pada tahun 2021. Meski menunjukkan tren penurunan, angka ini masih jauh dari target 14% yang diamanatkan Peraturan Presiden Nomor 72 Tahun 2021 tentang Percepatan Penurunan Stunting (Kementerian Kesehatan RI, 2022). Di Provinsi Jawa Tengah, prevalensi stunting berdasarkan SSGI 2022 tercatat sebesar 20,8% (Dinas Kesehatan Provinsi Jawa Tengah, 2023). Selain stunting, berbagai masalah gizi lain pada balita turut menjadi perhatian: gizi kurang/underweight sebesar 17,1%, gizi akut/wasting sebesar 7,7%, dan gizi lebih/overweight sebesar 4,5% berdasarkan SSGI 2022 (Kementerian Kesehatan RI, 2022). Kondisi multidimensi ini menunjukkan urgensi sistem klasifikasi status gizi berbasis data yang akurat dan efisien. Deteksi dini status gizi balita di tingkat fasilitas kesehatan pertama, seperti Puskesmas, sangat krusial untuk mencegah dampak buruk tersebut. Saat

ini, penentuan status gizi di Puskesmas Tirto, Kota Pekalongan, umumnya dilakukan berdasarkan pengukuran tubuh manusia (BB, TB, JK, dan Umur) yang kemudian dicocokkan dengan standar pertumbuhan. Namun, seiring dengan bertambahnya volume data balita, proses klasifikasi manual menjadi rentan terhadap kesalahan manusia (human error) dan menggunakan waktu yang lumayan lama. Karena alasan inilah, penerapan teknologi informasi melalui teknik Data Mining diperlukan untuk membantu tenaga kesehatan dalam mengklasifikasikan status gizi secara cepat dan akurat.

Data Mining atau penggalian data menawarkan solusi komputasi untuk menemukan pola tersembunyi dalam dataset besar. Dalam konteks klasifikasi kesehatan, algoritma Machine Learning sering digunakan untuk memprediksi kategori berdasarkan data latih. Dua algoritma klasifikasi yang populer dan memiliki karakteristik berbeda adalah K-Nearest Neighbor (KNN) dan Naive Bayes Classifier (NBC). KNN bekerja berdasarkan kedekatan jarak antar data (distance-based), sedangkan Naive Bayes bekerja berdasarkan perhitungan probabilitas statistik (probability-based). Pemilihan kedua algoritma ini didasari oleh kemampuannya dalam menangani data bertipe numerik dan kategorikal yang umum ditemukan pada data rekam medis balita.

Beberapa penelitian terdahulu telah menguji kinerja kedua algoritma ini dalam kasus serupa. Yuliansyah et al. [1] melakukan perbandingan metode KNN dan Naive Bayes untuk pengelompokan status gizi balita di Puskesmas, di mana ditemukan adanya perbedaan performa yang signifikan pada metrik presisi dan recall tergantung pada sebaran data. Selanjutnya, Bashir dan Aryani [2] membandingkan KNN, Naive Bayes, dan SVM di Puskesmas Sepatan, yang menyimpulkan bahwa akurasi model sangat dipengaruhi oleh tahap pra-pemrosesan data. Penelitian lain oleh Meriyana et al. [5] di Desa Pasirjengkol juga menerapkan kedua algoritma ini dengan atribut usia, jenis kelamin, dan tinggi badan, yang menunjukkan bahwa kedua metode tersebut cukup andal untuk deteksi dini stunting. Meskipun banyak studi telah dilakukan, kinerja algoritma Machine Learning sangat bergantung pada karakteristik dataset lokal yang digunakan. Perbedaan distribusi data demografis di Puskesmas Tirto, Kota Pekalongan, menuntut adanya pengujian spesifik untuk menentukan model mana yang paling optimal diterapkan di lokasi tersebut.

Berdasarkan uraian tersebut, penelitian ini bertujuan untuk membandingkan kinerja algoritma K-Nearest Neighbor (KNN) dan Naive Bayes Classifier (NBC) dalam mengklasifikasikan status gizi balita di Puskesmas Tirto, Kota Pekalongan. Penelitian ini menggunakan skema validasi hold-out dengan pembagian data 80% sebagai data latih (training) dan 20% sebagai data uji (testing) untuk mengukur kemampuan generalisasi model. Hasil komparasi ini diharapkan dapat memberikan rekomendasi algoritma terbaik yang dapat diimplementasikan sebagai alat bantu pendukung keputusan bagi tenaga medis dalam memantau status gizi balita secara lebih efisien.

2. TINJAUAN PUSTAKA

2.1. Studi Literatur

Penelitian mengenai komparasi algoritma klasifikasi untuk status gizi dan stunting balita telah banyak dilakukan dalam beberapa tahun terakhir. Namun, hasil yang diperoleh dari berbagai studi tersebut menunjukkan adanya variasi performa yang signifikan antara algoritma K-Nearest Neighbor (KNN) dan Naive Bayes Classifier (NBC), tergantung pada karakteristik dataset dan teknik pra-pemrosesan yang digunakan.

Beberapa penelitian terbaru menunjukkan keunggulan algoritma K-Nearest Neighbor (KNN). Studi yang dilakukan oleh Loka dan Marsal membandingkan kedua algoritma tersebut untuk klasifikasi status gizi balita, dan hasilnya menunjukkan bahwa KNN memperoleh akurasi sebesar 96,10%, lebih tinggi dibandingkan Naive Bayes yang hanya mencapai 90,94% [8]. Temuan ini didukung oleh penelitian Meriyana et al. di Desa Pasirjengkol yang menemukan dominasi mutlak KNN dengan akurasi 97,90% berbanding 54,39% untuk Naive Bayes [6]. Hal serupa juga disimpulkan oleh

Mu'taz dan Rosita pada prediksi stunting di Kecamatan Banjaran, di mana KNN dengan nilai $k = 3$ mencapai akurasi 99,11%, jauh mengungguli Naive Bayes yang terhambat oleh asumsi independensi fitur yang tidak terpenuhi pada data antropometri [2].

Di sisi lain, sejumlah penelitian justru menemukan bahwa Naive Bayes lebih unggul pada kondisi tertentu. Penelitian yang dilakukan oleh Alamin dan Aryani di Puskesmas Sepatan yang membandingkan KNN, Naive Bayes, dan Support Vector Machine (SVM) menyimpulkan bahwa Naive Bayes merupakan algoritma paling optimal dengan akurasi 99,65%, mengalahkan KNN yang memperoleh akurasi 96,36% [1]. Selain itu, Jannah et al. pada studi di Posyandu Cemara juga mencatat bahwa Naive Bayes memberikan hasil klasifikasi yang lebih optimal dalam mengidentifikasi kategori pertumbuhan anak dibandingkan KNN [5].

Perbedaan hasil yang kontradiktif dari penelitian-penelitian tersebut menunjukkan bahwa tidak ada satu algoritma yang secara universal terbaik untuk semua kasus data gizi balita. Performa model sangat dipengaruhi oleh distribusi data lokal, jumlah sampel, dan atribut yang digunakan. Oleh karena itu, penelitian ini bertujuan untuk menguji kinerja kedua algoritma tersebut secara spesifik pada dataset Puskesmas Tirto Kota Pekalongan dengan menggunakan skema validasi hold-out (80:20) guna menentukan model yang paling presisi untuk diterapkan di wilayah tersebut.

2.2. Data Mining

Data mining merupakan proses sistematis untuk memperoleh informasi bernilai yang sebelumnya tidak dapat diketahui secara langsung dari kumpulan data. Menurut Meriyana et al. [6], data mining melibatkan penggunaan teknik analisis statistik dan matematis yang didukung oleh kecerdasan buatan dan algoritma pembelajaran mesin untuk menghasilkan pengetahuan dan mengidentifikasi pola pengetahuan yang bermanfaat dari kumpulan data yang besar. Dalam sektor kesehatan, data mining menjadi alat vital untuk menganalisis rekam medis pasien guna menemukan pola penyakit atau kecenderungan status gizi yang tidak terlihat secara kasat mata.

2.3. Klasifikasi

Klasifikasi merupakan salah satu tugas utama dalam data mining. Alamin dan Aryani [1] menjelaskan bahwa klasifikasi adalah teknik pembelajaran mesin (machine learning) tipe supervised learning yang bertujuan untuk memetakan data input ke dalam kategori atau label kelas tertentu berdasarkan model yang dibentuk dari data pelatihan sebelumnya (training data). Proses ini bekerja dengan membangun sebuah model yang dapat memprediksi kelas dari data baru yang belum pernah dikenali sebelumnya. Akurasi klasifikasi sangat bergantung pada kualitas data latih dan pemilihan algoritma yang tepat untuk karakteristik data tersebut.

2.4. Status Gizi dan Stunting

Status gizi adalah penanda kesehatan yang menunjukkan balance antara asupan zat gizi dengan kebutuhan tubuh individu. Masalah gizi yang menjadi fokus global saat ini adalah stunting. Meriyana et al. [6] mengartikan stunting sebagai kondisi gangguan pertumbuhan pada kelompok usia balita akibat kekurangan gizi yang sangat amat parah, sehingga anak memiliki tinggi badan yang lebih pendek dibandingkan standar usianya. Penentuan status ini didasarkan pada indeks pengukuran ukuran tubuh manusia (BB, TB, JK, dan umur) yang kemudian digolongkan ke dalam kategori seperti gizi buruk, gizi kurang, gizi baik, atau stunting.

2.5. K-Nearest Neighbor(KNN)

K-Nearest Neighbor (KNN) adalah algoritma supervised learning yang sederhana namun efektif untuk klasifikasi. Prاتمanto et al. [4] menjelaskan bahwa prinsip kerja KNN adalah mengklasifikasikan data baru berdasarkan mayoritas kelas dari sejumlah k tetangga terdekatnya di dalam ruang fitur. Kedekatan ini dihitung menggunakan Euclidean distance. KNN memiliki keunggulan dalam menangani data dengan pola penyebaran yang tidak beraturan, namun kinerjanya sangat dipengaruhi oleh pemilihan nilai k dan skala data.

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

Dimana :

$d(x, y)$: nilai jarak antara vektor data uji (x) dan vektor data latih (y)

x_i : menyatakan atribut ke- i pada data test/uji

y_i : menyatakan atribut ke- i pada data training/latih

n : merupakan jumlah atribut atau dimensi data.

2.6. Naïve Bayes Classifier (NBC)

Naïve Bayes Classifier (NBC) adalah metode pengelompokan probabilistik berdasarkan pada Teorema Bayes. Karena penelitian ini menggunakan data atribut yang bersifat numerik (seperti BB, TB, umur, dan konversi JK), maka pendekatan yang digunakan adalah Gaussian Naive Bayes. Yuliansyah et al. [10] menjelaskan bahwa algoritma ini mengasumsikan nilai-nilai atribut mengikuti distribusi normal (Gaussian). Probabilitas setiap kelas dihitung menggunakan rumus densitas probabilitas Gaussian sebagai berikut:

$$P(x_i|c) = \frac{1}{\sqrt{2\pi\sigma_c^2}} * e^{-\frac{(x_i - \mu_c)^2}{2\sigma_c^2}} \quad (2)$$

Dimana :

$P(x_i|c)$: kemungkinan fitur x_i pada kelas c

μ_c : Rata-rata (mean) dari atribut pada kelas c

σ_c^2 : Varians dari atribut pada kelas c

x_i : Nilai atribut yang akan diklasifikasikan

π : Konstanta phi (3.14)

e : Bilangan Euler (2.718)

2.7. Hold-Out Method

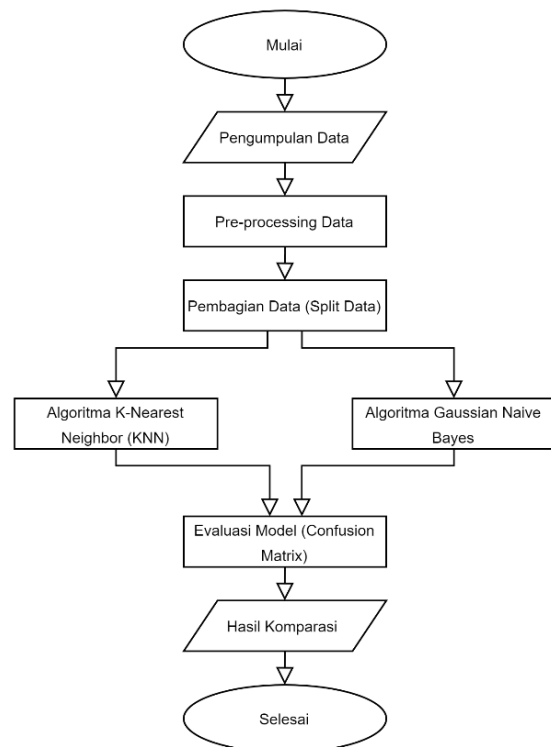
Evaluasi kinerja model merupakan tahapan krusial untuk memastikan bahwa algoritma yang dibangun mampu mengenali pola data dengan baik dan tidak hanya menghafal data latih (overfitting). Salah satu metode validasi yang umum digunakan adalah Hold-Out Method. Rais et al. [12] menjelaskan bahwa metode ini bekerja dengan cara membagi dataset menjadi dua kelompok, yaitu data latih sebagai dasar pembentukan model dan data uji untuk mengevaluasi performa model yang dihasilkan.

2.8. Python dan Google Colab

Untuk mengimplementasikan algoritma tersebut, penelitian ini menggunakan bahasa pemrograman Python. Mu'taz dan Rosita [2] menyebutkan bahwa Python adalah bahasa pemrograman yang sangat populer dalam bidang data science karena memiliki pustaka (library) yang lengkap seperti Pandas untuk manipulasi data, Scikit-Learn untuk algoritma machine learning, dan Matplotlib untuk visualisasi. Eksekusi kode dilakukan menggunakan Google Colab, sebuah lingkungan berbasis cloud yang memungkinkan pengolahan data tanpa memerlukan instalasi perangkat lunak yang berat di komputer lokal.

3. METODE PENELITIAN

Prosedur Penelitian Diagram alur penelitian disajikan pada Gambar 1 untuk menjelaskan tahapan apa saja yang dilakukan.



Gambar 1. Flowchart of the research procedure

3.1. Pengumpulan Data

Tahap awal dimulai dengan pengumpulan data sekunder berupa rekam medis balita dari Puskesmas Tirto, Kota Pekalongan. Dataset ini memuat atribut antropometri yang mencakup variabel usia, jenis kelamin, berat badan, tinggi badan, serta label status gizi.

3.2. Pre-processing

Tahap selanjutnya adalah pra-pemrosesan data, di mana data mentah diolah agar siap digunakan oleh algoritma machine learning. Proses ini diawali dengan pembersihan data (cleaning) yang bertujuan untuk memeriksa dan menangani keberadaan data kosong (missing values) atau data duplikat yang dapat mengganggu analisis. Setelah data bersih, dilakukan transformasi data (encoding) dengan mengubah atribut kategorikal, khususnya Jenis Kelamin (L/P), menjadi data numerik (0 atau 1) agar dapat diproses secara matematis oleh sistem. Langkah terakhir dalam tahap ini adalah normalisasi data menggunakan metode Min-Max Scaler untuk menyamakan skala nilai pada atribut numerik seperti Berat Badan dan Tinggi Badan, guna mencegah terjadinya bias pada perhitungan jarak dalam algoritma KNN.

3.3. Pembagian Data (Data Splitting)

Setelah melalui tahap pra-pemrosesan, dataset kemudian dibagi menggunakan metode validasi Hold-Out dengan rasio 80:20. Dalam skema ini, sebanyak 80% dari total data dialokasikan sebagai Data Latih (Training Set) yang berfungsi untuk melatih model dalam mengenali pola, sedangkan 20% sisanya dialokasikan sebagai Data Uji (Testing Set) yang digunakan khusus untuk mengevaluasi kinerja model yang telah terbentuk.

3.4. Implementasi Model Klasifikasi

Pada tahap implementasi, data latih diproses menggunakan dua algoritma klasifikasi yang berbeda untuk dibandingkan kinerjanya. Algoritma pertama adalah K-Nearest Neighbor (KNN), yang melakukan klasifikasi dengan cara menghitung kedekatan jarak (Euclidean Distance) antara data baru dengan data tetangga terdekatnya. Algoritma kedua adalah Gaussian Naive Bayes (NBC), yang melakukan klasifikasi berdasarkan perhitungan probabilitas

statistik dengan asumsi distribusi normal. Pendekatan Gaussian ini dipilih mengingat seluruh atribut data yang digunakan dalam penelitian ini telah dikonversi menjadi tipe data numerik.

3.5. Evaluasi Model

Model yang telah terbentuk diuji menggunakan Data Uji. Hasil prediksi dari kedua algoritma kemudian dianalisis menggunakan Confusion Matrix untuk menghitung nilai performa berupa Accuracy, Precision, Recall, dan F1-Score.

3.6. Analisis Hasil dan Kesimpulan

Tahap terakhir adalah membandingkan hasil performa antara KNN dan Naive Bayes untuk menyimpulkan algoritma mana yang memiliki kinerja terbaik dalam mengklasifikasikan status gizi balita di Puskesmas Tirto.

4. HASIL DAN PEMBAHASAN

4.1. Pre-processing Data

Tahap awal pengolahan data dimulai dengan menyeleksi atribut yang relevan dari dataset rekam medis balita yang berjumlah 203 data. Atribut yang digunakan untuk pemodelan meliputi Jenis Kelamin, Usia, Berat Badan, dan Tinggi Badan, dengan label target Status Gizi. Proses transformasi data yang dilakukan meliputi:

- Konversi Usia: Data usia diambil dari kolom 'Usia Saat Ukur' yang semula berformat teks (contoh: "2 Tahun - 0 Bulan - 8 Hari"). Nilai tahun dan bulan diekstraksi dari teks tersebut, kemudian dikonversi menjadi tipe data numerik dalam satuan bulan (usia_bulan) menggunakan rumus: (Tahun dikali 12) + Bulan.
- Encoding Jenis Kelamin: Atribut jenis kelamin (JK) diubah menjadi format numerik (JK_num), dimana Laki-laki (L) dikonversi menjadi 0 dan Perempuan (P) menjadi 1..
- Pembagian Data (Data Splitting) Dataset yang telah bersih dibagi dengan rasio 80:20. Dari total 203 data, sebanyak 162 data dialokasikan sebagai data latih dan 41 data sebagai data uji.

Berikut adalah sebelum dan sesudah dilakukan nya Proccesing data

Tabel 1. Pra-proccesing tabel

No	JK_num (0/1)	Usia	Berat (kg)	Tinggi (cm)	Status Gizi
1	L	2 Tahun - 0 Bulan - 8 Hari	11,7	89	Berat Badan Normal
2	P	1 Tahun - 3 Bulan - 7 Hari	8,3	75	Berat Badan Normal
3	L	1 Tahun - 2 Bulan - 8 Hari	9,0	75	Berat Badan Normal
4	P	2 Tahun - 0 Bulan - 4 Hari	9,0	82	Kurang
5	L	1 Tahun - 11 Bulan - 12 Hari	9,4	80	Kurang
.....					
203	P	1 Tahun - 7 Bulan - 2 Hari	8,5	78	Berat Badan Normal

Tabel 2. Pre-proccesing tabel

No	JK_num (0/1)	Usia (Bulan)	Berat (kg)	Tinggi (cm)	Status Gizi (Label)
1	0	24	11,7	89	Berat Badan Normal
2	1	15	8,3	75	Berat Badan Normal
3	0	14	9,0	75	Berat Badan Normal
4	1	24	9,0	82	Kurang
5	0	23	9,4	80	Kurang
.....					
203	1	19	8,5	78	Berat Badan Normal

4.2. Pengujian Algoritma K-Nearest Neighbor (KNN)

Pengujian algoritma K-Nearest Neighbor (KNN) dilakukan menggunakan data uji hasil pembagian data dengan metode hold-out, yaitu 80% data latih dan 20% data uji. Penentuan nilai parameter k dilakukan melalui proses eksperimen sistematis dengan menguji beberapa nilai k, yaitu $k = 3$, $k = 5$, $k = 7$, dan $k = 9$. Hasil pengujian menunjukkan bahwa nilai $k = 5$ menghasilkan akurasi tertinggi sebesar 85,37% pada data uji. Nilai $k = 3$ cenderung menghasilkan model yang sensitif terhadap noise (overfitting), sedangkan nilai k yang terlalu besar ($k = 7$ dan $k = 9$) menyebabkan batas keputusan menjadi terlalu umum (underfitting) karena tetangga yang diperhitungkan semakin jauh dan kurang representatif. Nilai $k = 5$ dipilih karena menawarkan keseimbangan optimal antara sensitivitas terhadap data lokal dan kemampuan generalisasi model. Pemilihan ini juga sesuai dengan rekomendasi umum dalam literatur bahwa nilai k ganjil yang moderat efektif untuk menghindari ambiguitas keputusan pada klasifikasi multi-kelas [4]. Model KNN dibangun menggunakan atribut umur (dalam bulan), berat badan, dan tinggi badan. Proses klasifikasi dilakukan dengan menghitung distance Euclidean data training dan data testing, kemudian menentukan kelas berdasarkan mayoritas kelas dari k tetangga terdekat. Hasil pengujian model KNN ditampilkan dalam bentuk confusion matrix pada Tabel berikut

Tabel 3. Confusion Matrix KNN

Prediksi	True Berat Normal	True Kurang	True Risiko	True Sangat Kurang
Normal	32	1	0	0
Kurang	3	3	0	0
Risiko Lebih	1	0	0	0
Sangat Kurang	0	1	0	0

Berdasarkan tabel tersebut, terlihat bahwa sebagian besar data dengan status gizi Normal berhasil diklasifikasikan dengan benar, yaitu sebanyak 32 data. Tetapi masih terdapat sedikit kesalahan pengelompokkan, terutama pada kelas Kurang dan Risiko Lebih yang sebagian diprediksi sebagai kelas Normal. Untuk mengukur kinerja model secara kuantitatif, digunakan metrik evaluasi akurasi, presisi,

recall, dan F1-score. Hasil evaluasi kinerja algoritma KNN ditunjukkan pada tabel berikut.

Tabel 4. Hasil Evaluasi Kinerja Algoritma KNN

Metrik	Nilai
Accuracy	85,37%
Precision	80,00%
Recall	85,00%
F1-Score	83,00%

Nilai *accuracy* sebesar 85,37% menunjukkan bahwa model KNN mampu mengklasifikasikan sebagian besar data uji dengan benar. Nilai *recall* yang cukup tinggi mengindikasikan bahwa model memiliki kemampuan yang baik dalam mengenali data dari setiap kelas status gizi. Selain itu, nilai *F1-score* sebesar 83,00% menunjukkan keseimbangan yang baik antara ketepatan prediksi dan kelengkapan hasil klasifikasi.

4.3. Pengujian Algoritma Naïve Bayes (NBC)

Pengujian algoritma Naive Bayes Classifier (NBC) dilakukan menggunakan data uji yang diperoleh dari pembagian data dengan metode hold-out, yaitu 80% sebagai data training sedangkan sisa data 20% sebagai data test. Algoritma ini bekerja berdasarkan pendekatan probabilistik dengan asumsi bahwa setiap atribut tidak memiliki ketergantungan antar variabel. Model Naive Bayes dibangun menggunakan atribut umur (bulan), berat badan, dan tinggi badan untuk memprediksi status gizi balita. Hasil pengujian model disajikan dalam bentuk confusion matrix pada tabel berikut.

Tabel 5. Confusion Matrix Naïve Bayes

Prediksi	True Berat Normal	True Kurang	True Risiko	True Sangat Kurang
Normal	26	1	0	6
Kurang	2	2	0	2
Risiko Lebih	1	1	0	0
Sangat Kurang	0	1	0	0

Berdasarkan hasil yang ditunjukkan pada Tabel 6, dapat dilihat bahwa algoritma Naive Bayes masih mengalami kesalahan klasifikasi yang cukup tinggi, terutama pada kelas Sangat Kurang dan Kurang yang sebagian besar diprediksi sebagai kelas Normal. Berdasarkan hasil tersebut, dapat disimpulkan bahwa

model mengalami kesulitan dalam membedakan kelas minoritas dengan jumlah data yang relatif sedikit. Untuk mengevaluasi performa model secara kuantitatif, digunakan metrik accuracy, precision, recall, dan F1-score. Hasil evaluasi kinerja algoritma Naive Bayes ditunjukkan pada tabel berikut

Tabel 6. Hasil Evaluasi Kinerja Algoritma KNN

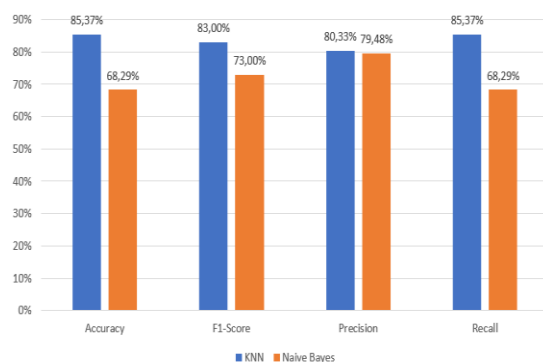
Metrik	Nilai
Accuracy	68,29%
Precision	79,00%
Recall	68,00 %
F1-Score	73,00%

Nilai *accuracy* sebesar 68,29% menunjukkan bahwa performa algoritma Naive Bayes masih berada di bawah algoritma KNN. Meskipun nilai *precision* relatif cukup tinggi, nilai *recall* yang lebih rendah mengindikasikan bahwa masih banyak data yang tidak berhasil dikenali dengan benar, khususnya pada kelas dengan jumlah data yang sedikit.

Nilai *F1-score* sebesar 73,00% menunjukkan bahwa keseimbangan antara ketepatan dan kelengkapan prediksi masih belum optimal. Hal ini dipengaruhi oleh asumsi independensi atribut pada algoritma Naive Bayes yang kurang sesuai dengan karakteristik data antropometri balita, di mana antar atribut saling berkorelasi

4.4. Komparasi Kinerja Algoritma

Setelah dilakukan pengujian terhadap algoritma K-Nearest Neighbor (KNN) dan Naive Bayes Classifier (NBC), selanjutnya dilakukan perbandingan kinerja kedua algoritma berdasarkan nilai accuracy, precision, recall, dan F1-score. Perbandingan ini bertujuan untuk mengetahui algoritma yang memberikan performa terbaik dalam klasifikasi status gizi balita. Hasil perbandingan kinerja kedua algoritma disajikan pada gambar berikut.



Gambar 2. Perbandingan Performa Algoritma

Berdasarkan hasil perbandingan tersebut, algoritma KNN menunjukkan kinerja yang secara konsisten lebih baik dibandingkan algoritma Naive Bayes pada hampir seluruh metrik evaluasi. Nilai accuracy KNN sebesar 85,37% jauh melampaui Naive

Bayes yang hanya mencapai 68,29%, dengan selisih sebesar 17,08 poin persentase.

Keunggulan KNN ini secara fundamental disebabkan oleh kesesuaian mekanisme kerjanya dengan karakteristik data yang digunakan. Data antropometri balita (umur, berat badan, tinggi badan) bersifat numerik kontinu dan memiliki korelasi antar atribut yang tinggi — misalnya, tinggi badan balita sangat dipengaruhi oleh usianya. KNN yang bekerja berdasarkan kedekatan jarak Euclidean mampu menangkap pola kemiripan lokal antar data secara langsung tanpa asumsi distribusi tertentu, sehingga lebih adaptif terhadap distribusi data nyata.

Sebaliknya, Naive Bayes mengandalkan asumsi independensi kondisional antar atribut (conditional independence assumption). Asumsi ini secara inheren melanggar pada data status gizi balita karena ketiga atribut utama — umur, berat badan, dan tinggi badan — memiliki korelasi yang kuat satu sama lain. Pelanggaran asumsi ini menyebabkan estimasi probabilitas posterior menjadi bias dan kurang akurat, khususnya untuk membedakan kelas-kelas yang memiliki rentang nilai atribut yang saling tumpang tindih seperti 'Kurang' dan 'Normal'.

Pada metrik recall, KNN memperoleh 85,37% dibanding Naive Bayes 68,29%. Perbedaan recall yang signifikan ini mengindikasikan bahwa Naive Bayes gagal mendeteksi banyak kasus positif sejati, terutama pada kelas minoritas seperti 'Sangat Kurang' dan 'Risiko Lebih'. Fenomena ini terjadi karena ketidakseimbangan kelas (class imbalance) dalam dataset — dominasi kelas 'Normal' membuat Naive Bayes cenderung memprediksi kelas mayoritas untuk meminimalkan error keseluruhan. KNN, dengan pendekatan berbasis tetangga lokal, lebih mampu menangani distribusi kelas yang tidak seimbang ini karena keputusan klasifikasi bergantung pada data di sekitarnya secara langsung.

Sementara itu, nilai precision kedua algoritma relatif tidak jauh berbeda, dengan KNN 80,33% dan Naive Bayes 79,48%. Hal ini dapat dijelaskan oleh fakta bahwa ketika Naive Bayes membuat prediksi positif, prediksi tersebut cukup dapat diandalkan, namun cakupannya jauh lebih terbatas (tercermin dari recall yang rendah). Perbedaan yang signifikan baru terlihat jelas pada F1-score — KNN 83,00% vs Naive Bayes 73,00% — yang merupakan rata-rata harmonik precision dan recall, sehingga lebih mencerminkan keseimbangan performa secara menyeluruh.

Berdasarkan analisis multidimensi tersebut, dapat disimpulkan bahwa algoritma K-Nearest Neighbor memiliki performa yang lebih stabil, lebih akurat, dan lebih andal dalam mengklasifikasikan status gizi balita pada dataset penelitian ini. Keunggulan KNN bukan semata karena kecanggihan algoritma, melainkan karena kesesuaian paradigma kerjanya dengan struktur dan karakteristik data antropometri yang bersifat numerik, berkorelasi tinggi, dan memiliki distribusi kelas yang tidak seimbang.

4.5. Pembahasan

Perbedaan performa antara algoritma K-Nearest Neighbor dan Naive Bayes menunjukkan bahwa karakteristik data memiliki peran penting dalam menentukan efektivitas metode penggolongan yang digunakan. Pada penelitian ini, atribut yang digunakan berupa umur, berat badan, dan tinggi badan, yang seluruhnya berbentuk data numerik dan saling berkaitan. KNN bekerja dengan membandingkan jarak antar data menggunakan ukuran kedekatan tertentu, sehingga data dengan karakteristik fisik yang mirip akan dikelompokkan dalam kelas yang sama. Pendekatan ini memungkinkan KNN menangkap pola alami yang terbentuk pada data status gizi balita, terutama ketika perbedaan antar kelas cukup dipengaruhi oleh variasi nilai antropometri. Sementara itu, Naive Bayes menghitung probabilitas kelas berdasarkan asumsi bahwa setiap atribut tidak saling memengaruhi. Pada data kesehatan seperti status gizi, asumsi ini sulit terpenuhi karena berat badan dan tinggi badan sangat dipengaruhi oleh umur. Ketergantungan antar atribut tersebut berpotensi menyebabkan perhitungan probabilitas menjadi kurang representatif terhadap kondisi sebenarnya. Selain faktor metode, jumlah data pada masing-masing kelas juga berpengaruh terhadap hasil klasifikasi. Dominasi data pada kategori berat badan normal menyebabkan model Naive Bayes cenderung lebih sering memprediksi kelas mayoritas, sehingga kemampuan mendeteksi kelas minoritas menjadi lebih rendah. Hal ini tercermin pada nilai *recall* yang lebih kecil dibandingkan KNN. Dengan mempertimbangkan mekanisme kerja algoritma dan struktur data yang digunakan, perbedaan performa yang diperoleh dalam penelitian ini dapat dipahami sebagai akibat dari kecocokan metode terhadap jenis data yang dianalisis, bukan semata-mata karena keunggulan satu algoritma secara umum.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian mengenai perbandingan antara algoritma K-Nearest Neighbor dengan Naive Bayes dalam pengelompokan status gizi balita menggunakan data dari Puskesmas Tirto, Kota Pekalongan, dapat disimpulkan bahwa proses klasifikasi dapat dilakukan dengan baik menggunakan pendekatan machine learning berbasis data antropometri, yaitu BB (berat badan), umur, TB (tinggi badan), dan jenis kelamin. Hasil pengujian menunjukkan bahwa algoritma K-Nearest Neighbor dengan nilai parameter $k = 5$ menghasilkan performa yang lebih bagus jika dibandingkan dengan algoritma Naive Bayes. Algoritma KNN memperoleh nilai akurasi mencapai 85,37%, presisi mencapai 80,33%, *recall* mencapai 85,37%, dan F1-score mencapai 83,00%. Sementara itu, algoritma Naive Bayes menghasilkan nilai akurasi mencapai 68,29%, presisi mencapai 79,48%, *recall* mencapai 68,29%, dan F1-score mencapai 73,00%. Perbedaan hasil tersebut menunjukkan bahwa algoritma KNN lebih sesuai digunakan pada data status gizi balita yang memiliki

atribut numerik dan saling berkaitan, dibandingkan dengan algoritma Naive Bayes yang mengandalkan asumsi independensi antar atribut. Dengan demikian, algoritma K-Nearest Neighbor dapat direkomendasikan sebagai metode yang lebih efektif untuk membantu proses klasifikasi status gizi balita pada dataset yang dipakai dalam penelitian ini.

DAFTAR PUSTAKA

- [1] P. B. Alamin and D. Aryani, "Perbandingan algoritma KNN, Naive Bayes, dan SVM untuk klasifikasi status stunting anak balita di Puskesmas Sepatan," *Jurnal Ilmiah Teknik dan Ilmu Komputer*, vol. 4, no. 4, pp. 360–369, 2025.
- [2] R. Fauzan Mu'taz and A. Rosita, "Comparison of KNN and Naive Bayes classification algorithms for predicting stunting in toddlers in Banjarn District," *Journal of Applied Informatics and Computing (JAIC)*, vol. 9, no. 5, pp. 2711–2717, 2025.
- [3] D. K. Harahap, E. T. Pardede, S. Nachwa, M. R. Maulizidan, and A. Ibrahim, "Komparasi kinerja KNN dan Naive Bayes untuk klasifikasi sentimen pengguna aplikasi DeepSeek AI," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 4, pp. 6021–6028, 2025.
- [4] Pratmanto, R. Wijianto, A. Widayanto, and Ubaidillah, "Komparasi K-Nearest Neighbors (KNN) dan Naive Bayes pada klasifikasi sentimen ulasan aplikasi Tokopedia di Google Play Store," *Informatics and Computer Engineering Journal*, vol. 5, no. 2, pp. 75–80, 2025.
- [5] M. Jannah, M. A. Hasan, and M. Al Fajar, "Perbandingan metode Naive Bayes dan K-Nearest Neighbor dalam mengklasifikasi status pertumbuhan anak stunting (studi kasus: Posyandu Cemara)," *Jurnal Fasilkom*, vol. 14, no. 1, pp. 250–255, 2024.
- [6] P. Meriyana, A. R. Pratama, E. Nurlaelasari, and A. R. Juwita, "Penerapan algoritma KNN dan Naive Bayes untuk klasifikasi stunting pada balita di Desa Pasirjengkol," *Jurnal Teknologi Informasi*, vol. 6, no. 2, pp. 156–168, 2025.
- [7] G. P. Insany, Somantri, and S. F. Maulina, "Penerapan algoritma Naive Bayes dan KNN pada klasifikasi gizi ibu hamil di Puskesmas Cicurug," *Jurnal Informatika Kesehatan*, vol. 9, pp. 710–719, 2024.
- [8] S. Kenia, P. Loka, and A. Marsal, "Comparison algorithm of K-Nearest Neighbor and Naive Bayes classifier for classifying nutritional status in toddlers," *Journal of Computer Science and Health Informatics*, vol. 3, pp. 8–14, 2023.
- [9] L. M. Sotarjua and D. B. Santoso, "Perbandingan algoritma KNN, Decision Tree, dan Random Forest pada data imbalanced class untuk klasifikasi promosi karyawan," *Jurnal Teknologi Informasi*, vol. 7, pp. 192–200, 2022.

- [10] M. R. Yuliansyah, B. Muslimin, and A. Franz, "Perbandingan metode K-Nearest Neighbors dan Naïve Bayes classifier pada klasifikasi status gizi balita di Puskesmas Muara Jawa Kota Samarinda," *Jurnal Informatika Terapan*, vol. 1, no. 1, pp. 8–20, 2022.
- [11] S. Alfaris and Kusnawi, "Komparasi metode KNN dan Naive Bayes terhadap analisis sentimen pengguna aplikasi," *Jurnal Sistem Informasi*, vol. 12, no. 1, pp. 2766–2776, 2023.
- [12] Z. Rais, M. K. Aidid, and A. D. Putri, "Penerapan algoritma K-Nearest Neighbor (K-NN) untuk analisis sentimen terhadap data ulasan aplikasi e-commerce Lazada pada Google Play Store," *Jurnal Variansi*, vol. 7, no. 2, pp. 143–154, 2025, doi: 10.35580/variantsium374.