



DOI: <https://doi.org/10.38035/dit.v4i1>
<https://creativecommons.org/licenses/by/4.0/>

Komparasi Algoritma K-Nearest Neighbor dan Naïve Bayes dalam Klasifikasi Status Gizi Balita Berbasis CRISP-DM

Almahira¹, Dony Oscar²

¹Universitas Bina Sarana Informatika, Jakarta, Indonesia, sofyanalmahira@gmail.com

²Universitas Bina Sarana Informatika, Jakarta, Indonesia, dony.dor@bsi.ac.id

Corresponding Author: sofyanalmahira@gmail.com¹

Abstract: Toddler nutritional status requires rapid and accurate classification so that children at risk can receive timely intervention. Manual matching of anthropometric records with reference tables at community health posts is time-consuming and vulnerable to recording and interpretation errors. This study compares K-Nearest Neighbor and Naïve Bayes for classifying toddler nutritional status using the Cross-Industry Standard Process for Data Mining framework. The initial dataset contained 232 anthropometric records collected in Pamanukan Kota Village in May 2026. After cleaning, 218 valid records were retained and divided into 174 training records and 44 testing records using an eighty-to-twenty stratified split. Predictors consisted of sex, age, body weight, and body height, while the target comprised four nutritional-status classes. Performance was evaluated using accuracy, precision, recall, F1-score, area under the curve, and computation time. K-Nearest Neighbor with one nearest neighbor achieved 79.55% accuracy, 57.81% precision, 59.19% recall, a 58.25% F1-score, and an area under the curve of 0.7418. Naïve Bayes achieved 47.73% accuracy and lower values on other predictive metrics. The novelty lies in an end-to-end comparison within a complete data-mining process on an imbalanced local dataset. K-Nearest Neighbor is recommended as the more adaptive model for developing a nutritional-status decision-support system.

Keywords: CRISP-DM, K-Nearest Neighbor, Machine Learning, Naïve Bayes, Toddler Nutritional Status

Abstrak: Klasifikasi status gizi balita perlu dilakukan secara cepat dan akurat agar anak yang berisiko dapat segera memperoleh intervensi. Pencocokan data antropometri dengan tabel rujukan secara manual di posyandu membutuhkan waktu dan rentan terhadap kesalahan pencatatan maupun interpretasi. Penelitian ini membandingkan kinerja K-Nearest Neighbor dan Naïve Bayes dalam mengklasifikasikan status gizi balita menggunakan kerangka Cross-Industry Standard Process for Data Mining. Dataset awal terdiri atas 232 rekam antropometri yang dikumpulkan di Posyandu Desa Pamanukan Kota pada Mei 2026. Setelah pembersihan data, diperoleh 218 rekam valid yang dibagi secara terstratifikasi menjadi 174 data pelatihan dan 44 data pengujian dengan rasio delapan puluh berbanding dua puluh. Variabel prediktor meliputi jenis kelamin, umur, berat badan, dan tinggi badan, sedangkan target terdiri atas empat kelas status gizi. Kinerja model dievaluasi menggunakan akurasi, presisi, recall, F1-score, area under the curve, dan waktu komputasi. K-Nearest Neighbor dengan satu tetangga terdekat

menghasilkan akurasi 79,55%, presisi 57,81%, recall 59,19%, F1-score 58,25%, dan area under the curve 0,7418. Naïve Bayes menghasilkan akurasi 47,73% serta nilai metrik prediktif lain yang lebih rendah. Kebaruan penelitian terletak pada komparasi menyeluruh dalam kerangka proses penambangan data lengkap pada dataset lokal posyandu yang tidak seimbang. K-Nearest Neighbor direkomendasikan untuk dikembangkan menjadi sistem pendukung keputusan status gizi.

Kata Kunci: CRISP-DM, K-Nearest Neighbor, Machine Learning, Naïve Bayes, Status Gizi Balita

PENDAHULUAN

Kesehatan dan pertumbuhan balita merupakan indikator penting bagi kualitas sumber daya manusia karena gangguan gizi pada periode awal kehidupan dapat memengaruhi perkembangan fisik dan kognitif dalam jangka panjang. Standar pertumbuhan anak menggunakan indikator antropometri, antara lain berat badan menurut umur, tinggi badan menurut umur, dan berat badan menurut tinggi badan, untuk menilai apakah pertumbuhan seorang anak berada pada rentang yang sesuai (World Health Organization, 2026). Di Indonesia, penentuan kategori status gizi anak mengacu pada Peraturan Menteri Kesehatan Republik Indonesia Nomor 2 Tahun 2020 tentang Standar Antropometri Anak (Kementerian Kesehatan Republik Indonesia, 2020).

Masalah malnutrisi masih membutuhkan perhatian serius. Estimasi bersama United Nations Children's Fund, World Health Organization, dan World Bank menunjukkan bahwa pada 2024 terdapat 150,2 juta anak balita yang mengalami stunting, 42,8 juta mengalami wasting, dan 35,5 juta mengalami kelebihan berat badan secara global (World Health Organization, 2025). Di Indonesia, prevalensi stunting dilaporkan turun dari 24,4% pada 2021 menjadi 21,6% pada 2022 (Kementerian Kesehatan Republik Indonesia, 2023), kemudian menjadi 19,8% berdasarkan Survei Status Gizi Indonesia 2024 (Kementerian Kesehatan Republik Indonesia, 2024). Penurunan tersebut merupakan perkembangan positif, tetapi ketepatan identifikasi anak dengan gizi buruk, gizi kurang, gizi baik, atau gizi lebih tetap menentukan kecepatan intervensi di tingkat layanan dasar.

Pos Pelayanan Terpadu menjadi salah satu titik terdekat bagi masyarakat untuk memantau pertumbuhan balita. Namun, observasi pada Posyandu Desa Pamanukan Kota menunjukkan bahwa data antropometri yang telah dikumpulkan masih dicocokkan secara manual dengan tabel standar. Proses tersebut membutuhkan waktu, rentan terhadap kesalahan pencatatan dan interpretasi, serta menyebabkan data yang telah tersedia belum sepenuhnya diolah menjadi informasi pendukung keputusan. Permasalahan serupa juga ditemukan dalam pengembangan sistem klasifikasi status gizi berbasis web, yang menunjukkan perlunya otomatisasi untuk mengurangi *human error* dan mempercepat pengolahan data (Dwiyanti and Saptari, 2023).

Pemanfaatan *machine learning* dapat membantu mengubah rekam antropometri menjadi hasil klasifikasi yang lebih cepat dan konsisten. K-Nearest Neighbor mengelompokkan data baru berdasarkan kedekatan dengan data pelatihan, sedangkan Naïve Bayes menggunakan probabilitas kelas berdasarkan karakteristik atribut (Han, Kamber and Pei, 2011). Kedua algoritma tersebut relatif sederhana dan banyak digunakan pada dataset kesehatan. Akan tetapi, performanya tidak selalu konsisten. (Loka and Marsal, 2023) menemukan K-Nearest Neighbor lebih unggul dengan akurasi 96,24% dibandingkan Naïve Bayes sebesar 91,00%, sedangkan (Pangumbara'an, Wakista and Makhsun, 2024) melaporkan Naïve Bayes mencapai 87,65% dan K-Nearest Neighbor 76,19%. Perbedaan tersebut menegaskan bahwa karakteristik dataset, distribusi kelas, skala atribut, dan tahapan prapemrosesan dapat memengaruhi hasil model.

Penelitian lain menunjukkan bahwa komparasi algoritma pada data stunting dan status gizi terus berkembang. (Dwiyanti and Saptari, 2023) membandingkan K-Nearest Neighbor, Naïve Bayes, dan Support Vector Machine untuk klasifikasi risiko stunting, sedangkan (Hardiani and Putri, 2024) menerapkan Naïve Bayes dalam kerangka klasifikasi stunting. Tinjauan sistematis oleh (Indrisari, Febiansyah and Adiwinoto, 2025) mengidentifikasi ketidakseimbangan kelas, keterbatasan interpretabilitas, serta minimnya validasi eksternal sebagai tantangan utama penerapan *machine learning* untuk prediksi stunting di Indonesia. Meskipun demikian, komparasi khusus K-Nearest Neighbor dan Naïve Bayes yang mengikuti yang tidak seimbang masih terbatas.

Sejumlah penelitian lain turut memperlihatkan beragam kinerja algoritma dalam klasifikasi status gizi balita. (Kotler, Kartajaya and Setiawan, 2021) mengembangkan klasifikasi berbasis web menggunakan Naïve Bayes dan K-Nearest Neighbor, sedangkan Harliana et al. (2022) menggabungkan Naïve Bayes dengan pemetaan GIS untuk mendukung monitoring status gizi. (Lonang and Normawati, 2022) menunjukkan bahwa seleksi fitur Backward Elimination dapat meningkatkan kinerja K-Nearest Neighbor pada klasifikasi stunting. Pengujian Naïve Bayes dengan K-Fold Cross Validation juga dilakukan oleh (Nurainun et al., 2023), sementara (Kurniawan, Rahim and Siswa, 2024) menerapkan Gaussian Naïve Bayes pada klasifikasi status gizi balita. Pada penelitian yang menggunakan alur CRISP-DM, (Islam, Mulyadien and Enri, 2022) menguji algoritma C4.5, sedangkan (Kurniawan, Rahim and Siswa, 2024) menerapkan Naive Bayes berdasarkan pengukuran antropometri. Beragam hasil tersebut memperkuat bahwa pemilihan model dan rancangan evaluasi perlu disesuaikan dengan karakteristik data yang digunakan.

Research gap penelitian ini terletak pada belum banyaknya pengujian dua algoritma tersebut secara mendalam melalui alur yang mencakup pemahaman kebutuhan, eksplorasi data, pembersihan dan transformasi, pemodelan, evaluasi, serta rekomendasi implementasi pada satu dataset lokal. Keunikan atau *state of the art* penelitian terletak pada penggunaan kerangka Cross-Industry Standard Process for Data Mining secara menyeluruh untuk membandingkan performa K-Nearest Neighbor dan Naïve Bayes pada empat kelas status gizi dengan distribusi yang tidak seimbang. Evaluasi tidak hanya menggunakan akurasi, tetapi juga presisi, *recall*, *F1-score*, *area under the curve*, dan waktu komputasi sehingga pemilihan model tidak bergantung pada satu ukuran.

Berdasarkan uraian tersebut, penelitian ini bertujuan menerapkan tahapan Cross-Industry Standard Process for Data Mining pada data antropometri balita, membandingkan performa K-Nearest Neighbor dan Naïve Bayes, serta menentukan algoritma yang paling adaptif untuk dataset Posyandu Desa Pamanukan Kota. Hasil penelitian diharapkan menjadi dasar pengembangan sistem pendukung keputusan yang membantu kader posyandu mengklasifikasikan status gizi secara lebih cepat, objektif, dan terukur.

METODE

Penelitian menggunakan pendekatan kuantitatif komparatif dengan kerangka Cross-Industry Standard Process for Data Mining. Kerangka tersebut terdiri atas enam fase, yaitu *business understanding*, *data understanding*, *data preparation*, *modeling*, *evaluation*, dan *deployment* (Chapman et al., 2000). Objek penelitian berupa rekam antropometri balita dari Posyandu Desa Pamanukan Kota yang dikumpulkan pada Mei 2026. Dataset awal berjumlah 232 rekam. Setelah pemeriksaan kelengkapan dan kewajaran nilai, 218 rekam dinyatakan valid dan digunakan pada tahap pemodelan.

Variabel prediktor terdiri atas jenis kelamin, umur dalam bulan, berat badan dalam kilogram, serta tinggi atau panjang badan dalam sentimeter. Nama balita dikeluarkan karena tidak memiliki kemampuan prediktif, sedangkan *Z-score* tidak digunakan sebagai prediktor untuk mencegah *data leakage* karena nilai tersebut menjadi dasar pembentukan label. Variabel target adalah status gizi yang terdiri atas Gizi Buruk, Gizi Kurang, Gizi Baik, dan Gizi Lebih

sesuai standar antropometri anak (Kementerian Kesehatan Republik Indonesia, 2020). Distribusi data bersih menunjukkan dominasi kelas Gizi Baik sebanyak 135 rekam, sedangkan kelas lain memiliki jumlah lebih kecil.

Tahap *data preparation* mencakup pemilihan atribut, pembersihan data kosong atau tidak wajar, pengkodean variabel kategorik, normalisasi, dan pembagian dataset. Jenis kelamin dikodekan menjadi nilai numerik. Atribut numerik untuk K-Nearest Neighbor dinormalisasi menggunakan Min-Max Scaler pada rentang 0 sampai 1 karena algoritma berbasis jarak sensitif terhadap perbedaan skala. Naïve Bayes dijalankan pada data tanpa normalisasi. Dataset dibagi menggunakan *stratified hold-out* 80:20 sehingga diperoleh 174 data pelatihan dan 44 data pengujian. Pembagian terstratifikasi digunakan untuk mempertahankan proporsi setiap kelas.

Pada tahap *modeling*, K-Nearest Neighbor diuji menggunakan nilai K ganjil dari 1 sampai 21 dan jarak Euclidean. Nilai K terbaik dipilih berdasarkan akurasi tertinggi pada data pengujian. Model pembandingan menggunakan Gaussian Naïve Bayes karena atribut prediktor bersifat numerik kontinu. Seluruh eksperimen dilakukan pada Google Colaboratory menggunakan Python dan pustaka Scikit-learn, Pandas, NumPy, dan Matplotlib.

Tahap *evaluation* menggunakan *confusion matrix* untuk memperoleh akurasi, presisi, *recall*, *F1-score*, dan *area under the curve*. Akurasi menunjukkan proporsi prediksi benar secara keseluruhan; presisi mengukur ketepatan prediksi suatu kelas; *recall* mengukur kemampuan menemukan anggota kelas aktual; sedangkan *F1-score* menyeimbangkan presisi dan *recall*. Nilai makro digunakan untuk menggambarkan performa lintas kelas pada dataset yang tidak seimbang. Waktu pelatihan dan pengujian turut dibandingkan. Pada tahap *deployment*, algoritma dengan kinerja prediktif terbaik direkomendasikan sebagai dasar sistem pendukung keputusan, tetapi belum diterapkan sebagai aplikasi operasional.

Tabel 1. Distribusi Kelas dan Pembagian Dataset

Status Gizi	Data Bersih	Pelatihan	Pengujian	Persentase
Gizi Baik	135	108	27	61,93%
Gizi Kurang	36	29	7	16,51%
Gizi Buruk	27	21	6	12,39%
Gizi Lebih	20	16	4	9,17%
Total	218	174	44	100,00%

Sumber: Data primer yang diolah, 2026

HASIL DAN PEMBAHASAN

Pemahaman dan Persiapan Data

Tahap *data understanding* menunjukkan bahwa dataset yang telah melalui proses seleksi awal terdiri atas 218 rekam valid dengan distribusi kelas yang tidak seimbang. Kelas Gizi Baik menjadi kelas mayoritas dengan jumlah 135 rekam atau 61,93% dari keseluruhan data. Sementara itu, kelas Gizi Kurang berjumlah 36 rekam, kelas Gizi Buruk sebanyak 27 rekam, dan kelas Gizi Lebih hanya berjumlah 20 rekam atau 9,17%. Distribusi setiap kelas tersebut disajikan pada Gambar 1. Komposisi ini memperlihatkan bahwa sebagian besar balita dalam dataset berada pada kategori Gizi Baik, sedangkan jumlah data pada kategori lainnya relatif lebih sedikit, khususnya kelas Gizi Lebih sebagai kelas minoritas.

Ketidakseimbangan distribusi kelas menjadi salah satu karakteristik penting yang perlu diperhatikan dalam proses pemodelan. Dominasi kelas Gizi Baik dapat menyebabkan model lebih sering mempelajari pola kelas mayoritas sehingga cenderung menghasilkan prediksi ke arah kategori tersebut. Kondisi ini berpotensi membuat model memiliki akurasi keseluruhan yang tinggi, tetapi belum tentu mampu mengidentifikasi kelas Gizi Buruk, Gizi Kurang, dan Gizi Lebih secara optimal. Dalam konteks pemantauan status gizi balita, kesalahan dalam mengenali kelas minoritas perlu diperhatikan karena setiap kategori memiliki kebutuhan penanganan dan tindak lanjut yang berbeda.

Oleh karena itu, evaluasi kinerja model tidak cukup dilakukan hanya berdasarkan nilai akurasi. Akurasi hanya menunjukkan proporsi keseluruhan prediksi yang benar tanpa menjelaskan kemampuan model dalam mengenali setiap kategori status gizi. Evaluasi perlu dilengkapi dengan nilai presisi, *recall*, dan *F1-score* pada masing-masing kelas. Presisi digunakan untuk mengetahui ketepatan prediksi model pada suatu kategori, sedangkan *recall* menunjukkan kemampuan model dalam menemukan seluruh data aktual pada kategori tersebut. Adapun *F1-score* memberikan gambaran keseimbangan antara presisi dan *recall*. Penggunaan beberapa metrik tersebut memungkinkan penilaian performa model dilakukan secara lebih menyeluruh, terutama pada dataset yang memiliki distribusi kelas tidak seimbang.

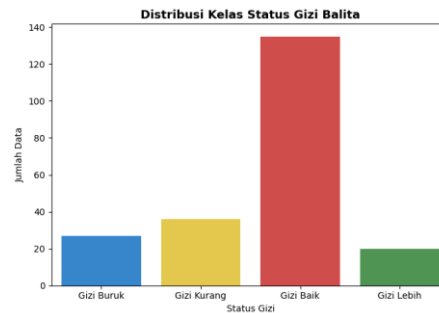
Tahap berikutnya adalah *data preparation* yang dilakukan untuk menghasilkan dataset yang lebih bersih, konsisten, dan sesuai dengan kebutuhan algoritma. Proses pembersihan data menghasilkan 218 rekam yang memenuhi kriteria untuk digunakan dalam pemodelan. Data yang tidak lengkap, tidak sesuai dengan batas usia balita, atau memiliki nilai antropometri yang tidak wajar dikeluarkan dari dataset. Langkah tersebut dilakukan agar model tidak mempelajari pola yang berasal dari kesalahan pencatatan atau nilai ekstrem yang tidak mencerminkan kondisi sebenarnya.

Dalam tahap seleksi atribut, nama atau identitas balita tidak digunakan karena tidak memiliki hubungan langsung dengan proses klasifikasi status gizi. Nilai *Z-score* berat badan menurut umur dan tinggi badan menurut umur juga tidak dimasukkan sebagai variabel prediktor. Keputusan mengecualikan *Z-score* merupakan langkah penting untuk mencegah terjadinya *data leakage*, yaitu kondisi ketika model memperoleh informasi yang secara langsung berkaitan dengan pembentukan label kelas. Apabila *Z-score* digunakan sebagai prediktor, model berpotensi menghasilkan kinerja yang tampak sangat tinggi, tetapi hasil tersebut tidak sepenuhnya menggambarkan kemampuan algoritma dalam mempelajari hubungan antara karakteristik antropometri dan status gizi.

Variabel yang dipertahankan sebagai masukan model meliputi jenis kelamin, umur, berat badan, dan tinggi badan. Keempat variabel tersebut digunakan untuk membangun pola klasifikasi status gizi melalui algoritma K-Nearest Neighbor dan Naïve Bayes. Variabel jenis kelamin terlebih dahulu dikodekan ke dalam bentuk numerik agar dapat diproses oleh algoritma, sedangkan variabel umur, berat badan, dan tinggi badan dipertahankan sebagai data numerik.

Normalisasi data diterapkan secara khusus pada model K-Nearest Neighbor karena algoritma tersebut melakukan klasifikasi berdasarkan perhitungan jarak antardata. Perbedaan rentang nilai antara umur, berat badan, dan tinggi badan dapat menyebabkan variabel yang memiliki skala lebih besar memberikan pengaruh lebih dominan dalam perhitungan jarak. Melalui normalisasi, seluruh variabel ditempatkan pada skala yang sebanding sehingga kontribusi setiap atribut menjadi lebih proporsional. Sementara itu, Naïve Bayes tidak memerlukan proses normalisasi yang sama karena proses klasifikasinya didasarkan pada perhitungan probabilitas.

Setelah seluruh tahapan pembersihan, pengkodean, dan transformasi selesai dilakukan, dataset dibagi menjadi data pelatihan dan data pengujian menggunakan rasio 80:20. Pembagian tersebut menghasilkan 174 data pelatihan dan 44 data pengujian. Proses pembagian dilakukan secara terstratifikasi agar proporsi setiap kelas pada data pelatihan dan data pengujian tetap mencerminkan distribusi kelas pada dataset bersih. Dengan demikian, seluruh kategori status gizi tetap terwakili dalam kedua bagian data, sehingga proses pelatihan dan evaluasi model dapat dilakukan secara lebih konsisten.



Gambar 1. Distribusi Kelas Status Gizi Balita
Sumber: Output Google Colab, 2026

Hasil Pemodelan K-Nearest Neighbor

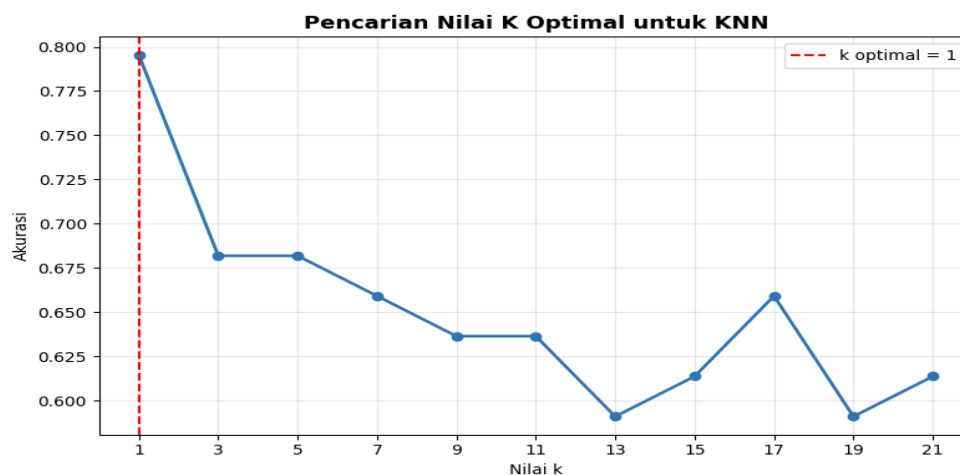
Eksperimen nilai K menunjukkan bahwa $K = 1$ menghasilkan akurasi tertinggi, yaitu 79,55%. Akurasi menurun menjadi 68,18% pada $K = 3$ dan $K = 5$, kemudian berfluktuasi pada kisaran 59,09% sampai 65,91% untuk nilai K yang lebih besar. Pola tersebut menunjukkan bahwa tetangga terdekat tunggal paling sesuai dengan struktur data lokal yang digunakan. Ketika nilai K bertambah, keputusan klasifikasi semakin dipengaruhi oleh kelas mayoritas sehingga batas antar kelas minoritas menjadi kurang jelas. Hasil pencarian nilai K disajikan pada Tabel 2 dan Gambar 2.

Pemilihan $K = 1$ juga menunjukkan bahwa data antropometri memiliki kelompok lokal yang cukup spesifik. Namun, nilai K yang sangat kecil dapat sensitif terhadap penciran dan variasi sampel. Karena itu, keunggulan $K = 1$ pada penelitian ini perlu dibaca sebagai hasil terbaik pada dataset dan pembagian data yang digunakan, bukan sebagai nilai universal untuk seluruh data posyandu.

Tabel 2. Hasil Pencarian Nilai K Optimal

Nilai K	Akurasi	Nilai K	Akurasi
1	79,55%	3	68,18%
5	68,18%	7	65,91%
9	63,64%	11	63,64%
13	59,09%	15	61,36%
17	65,91%	19	59,09%
21	61,36%		

Sumber: Output pengujian yang diolah, 2026



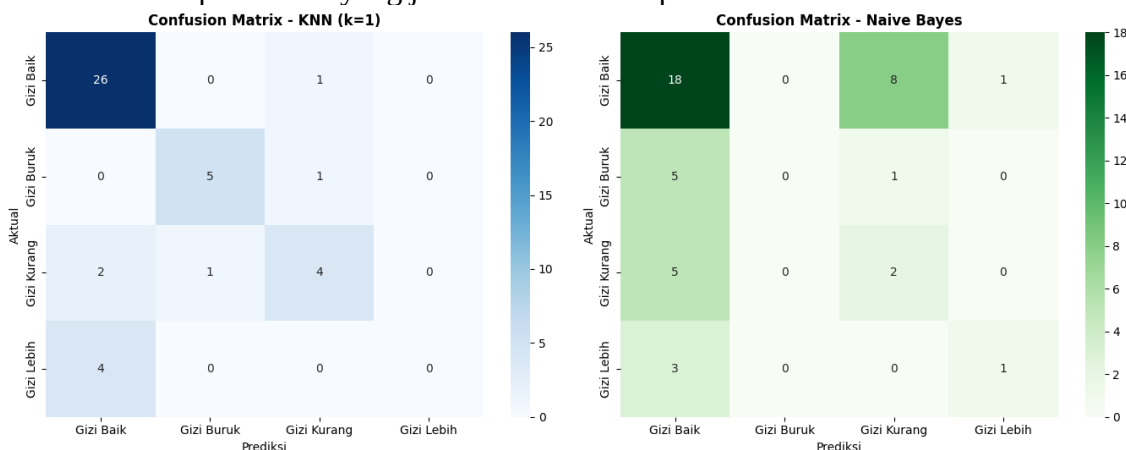
Gambar 2. Pencarian Nilai K Optimal untuk K-Nearest Neighbor
Sumber: Output Google Colab, 2026

Evaluasi K-Nearest Neighbor dan Naïve Bayes

Model K-Nearest Neighbor dengan $K = 1$ menghasilkan 35 prediksi benar dari 44 data pengujian sehingga memperoleh akurasi 79,55%. Nilai presisi makro mencapai 57,81%, *recall* 59,19%, *F1-score* 58,25%, dan *area under the curve* 0,7418. Pada tingkat kelas, performa terbaik ditemukan pada Gizi Baik dengan *F1-score* 0,88 dan Gizi Buruk sebesar 0,83. Kinerja pada Gizi Kurang masih moderat dengan *F1-score* 0,62, sedangkan seluruh empat data Gizi Lebih salah diprediksi sehingga *F1-score* kelas tersebut bernilai 0,00.

Naïve Bayes menghasilkan 21 prediksi benar dari 44 data pengujian dengan akurasi 47,73%. Presisi makro sebesar 31,56%, *recall* 30,06%, *F1-score* 29,41%, dan *area under the curve* 0,6914. Model masih mampu mengenali sebagian besar data Gizi Baik, tetapi tidak berhasil mengidentifikasi kelas Gizi Buruk pada data pengujian. Kinerja ini menunjukkan bahwa asumsi independensi antar atribut pada Naïve Bayes belum sesuai sepenuhnya dengan hubungan biologis antara umur, berat badan, dan tinggi badan. Selain itu, probabilitas prior yang lebih tinggi pada kelas Gizi Baik mendorong prediksi menuju kelas mayoritas.

Perbandingan *confusion matrix* pada Gambar 3 memperlihatkan perbedaan pola kesalahan kedua model. K-Nearest Neighbor relatif mampu mempertahankan prediksi pada tiga kelas, meskipun gagal mengenali Gizi Lebih. Naïve Bayes menghasilkan kesalahan yang lebih tersebar dan tidak memberikan prediksi benar untuk Gizi Buruk. Temuan tersebut menguatkan pentingnya membaca performa per kelas ketika dataset tidak seimbang. Akurasi yang tampak cukup tinggi pada kelas mayoritas belum menjamin kemampuan model mendeteksi balita pada kelas yang justru memerlukan perhatian lebih besar.



Gambar 3. Confusion Matrix K-Nearest Neighbor dan Naïve Bayes

Sumber: Output Google Colab, 2026

Tabel 3. Hasil Evaluasi Model per Kelas

Model/Kelas	Presisi	Recall	F1-Score	Support
KNN - Gizi Baik	0,81	0,96	0,88	27
KNN - Gizi Buruk	0,83	0,83	0,83	6
KNN - Gizi Kurang	0,67	0,57	0,62	7
KNN - Gizi Lebih	0,00	0,00	0,00	4
Naïve Bayes - Gizi Baik	0,58	0,67	0,62	27
Naïve Bayes - Gizi Buruk	0,00	0,00	0,00	6
Naïve Bayes - Gizi Kurang	0,18	0,29	0,22	7
Naïve Bayes - Gizi Lebih	0,50	0,25	0,33	4

Sumber: Output pengujian yang diolah, 2026

Komparasi Kinerja dan Pembahasan

Secara keseluruhan, K-Nearest Neighbor unggul pada seluruh metrik prediktif. Selisih akurasi antara kedua model mencapai 31,82 poin persentase. Selisih presisi, *recall*, dan F1-score masing-masing sebesar 26,25, 29,13, dan 28,84 poin persentase. Nilai *area under the curve* K-Nearest Neighbor juga lebih tinggi sebesar 0,0504. Waktu pelatihan K-Nearest Neighbor lebih cepat, yaitu 3,92 milidetik dibandingkan 7,72 milidetik pada Naïve Bayes. Naïve Bayes hanya lebih cepat pada waktu pengujian, yaitu 2,55 milidetik dibandingkan 4,47 milidetik, tetapi perbedaan beberapa milidetik tersebut tidak sebanding dengan penurunan kemampuan klasifikasi.

Hasil ini sejalan dengan (Loka and Marsal, 2023), yang juga menemukan K-Nearest Neighbor lebih akurat daripada Naïve Bayes untuk klasifikasi status gizi balita. Kedekatan pola umur, berat badan, dan tinggi badan memungkinkan algoritma berbasis jarak mengenali kelompok data yang serupa secara lokal. Sebaliknya, hasil penelitian ini berbeda dari (Pangumbara'an, Wakista and Makhsun, 2024), yang melaporkan Naïve Bayes lebih unggul. Perbedaan tersebut dapat dipengaruhi oleh jumlah sampel, komposisi kelas, fitur yang digunakan, proses normalisasi, serta skema pembagian data. Dengan demikian, pemilihan algoritma perlu didasarkan pada evaluasi terhadap dataset sasaran, bukan semata-mata pada performa yang dilaporkan penelitian sebelumnya.

Meskipun unggul, K-Nearest Neighbor belum memberikan hasil memadai pada kelas Gizi Lebih. Kegagalan mengenali empat data uji pada kelas tersebut berkaitan dengan jumlah data yang paling sedikit dan kemungkinan kedekatan karakteristiknya dengan kelas lain. Temuan ini konsisten dengan (Indrisari, Febiansyah and Adiwino, 2025), yang menempatkan *imbalanced data* sebagai salah satu masalah utama dalam pengembangan model prediksi stunting di Indonesia. Oleh karena itu, akurasi 79,55% tidak boleh ditafsirkan bahwa model telah siap digunakan sebagai alat diagnosis. Model lebih tepat diposisikan sebagai alat bantu penyaringan awal yang tetap memerlukan verifikasi kader atau tenaga kesehatan.

Tabel 4 dan Gambar 4 menegaskan bahwa rekomendasi K-Nearest Neighbor didasarkan pada keseimbangan performa, bukan hanya akurasi. Model mampu menghasilkan F1-score hampir dua kali lipat dibandingkan Naïve Bayes dan memiliki diskriminasi kelas yang lebih baik berdasarkan *area under the curve*. Dalam konteks operasional posyandu, tambahan waktu pengujian sekitar dua milidetik tidak menjadi hambatan berarti karena jumlah data yang diproses relatif terbatas.

Kebaruan penelitian ini terletak pada bukti empiris bahwa penerapan alur Cross-Industry Standard Process for Data Mining secara lengkap pada dataset lokal posyandu yang tidak seimbang menghasilkan pemilihan model yang berbeda dari sebagian penelitian terdahulu. Pengujian per kelas menunjukkan bahwa model terbaik secara agregat masih dapat gagal pada kelas minoritas tertentu. Temuan tersebut memperluas pembahasan komparasi algoritma dengan menekankan bahwa rekomendasi implementasi harus mempertimbangkan distribusi kelas dan metrik makro. Berdasarkan hasil ini, K-Nearest Neighbor dengan $K = 1$ direkomendasikan sebagai model awal untuk pengembangan sistem pendukung keputusan status gizi di Posyandu Desa Pamanukan Kota, disertai pengayaan data dan validasi lebih lanjut sebelum digunakan secara operasional.

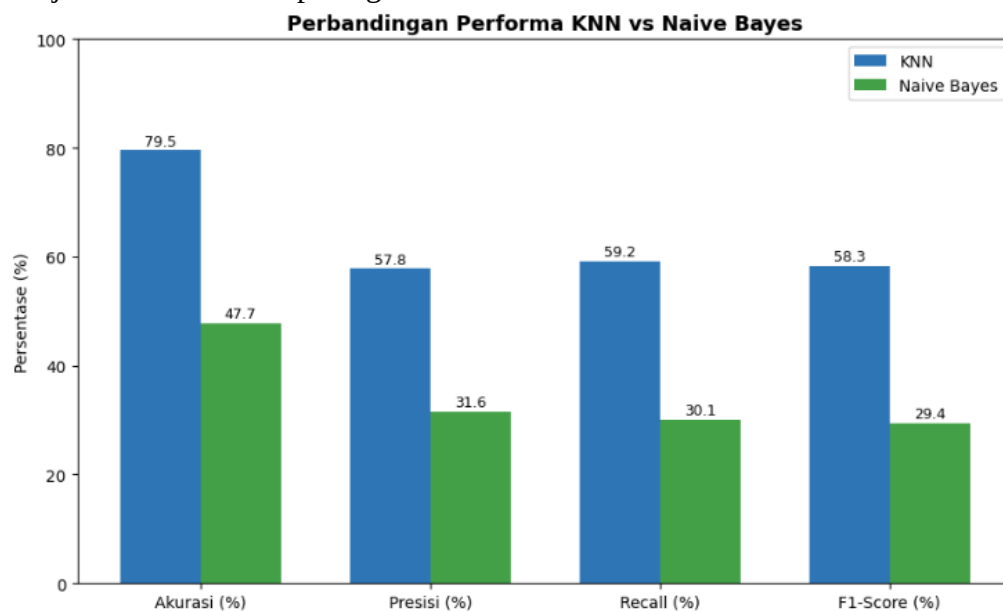
Tabel 4. Komparasi Performa K-Nearest Neighbor dan Naïve Bayes

Metrik	K-Nearest Neighbor	Naïve Bayes	Model Terbaik
Akurasi	79,55%	47,73%	K-Nearest Neighbor
Presisi Makro	57,81%	31,56%	K-Nearest Neighbor
Recall Makro	59,19%	30,06%	K-Nearest Neighbor
F1-Score Makro	58,25%	29,41%	K-Nearest Neighbor

Metrik	K-Nearest Neighbor	Naïve Bayes	Model Terbaik
Area Under the Curve	0,7418	0,6914	K-Nearest Neighbor
Waktu Pelatihan	3,92 ms	7,72 ms	K-Nearest Neighbor
Waktu Pengujian	4,47 ms	2,55 ms	Naïve Bayes

Sumber: Output pengujian yang diolah, 2026

Berdasarkan Tabel 4, K-Nearest Neighbor menunjukkan kinerja yang lebih baik daripada Naïve Bayes pada hampir seluruh metrik evaluasi, yaitu akurasi, presisi makro, *recall* makro, *F1-score* makro, dan *area under the curve*. K-Nearest Neighbor juga memiliki waktu pelatihan yang lebih singkat, sedangkan Naïve Bayes hanya lebih unggul dalam waktu pengujian. Hasil tersebut menunjukkan bahwa K-Nearest Neighbor merupakan model terbaik secara keseluruhan untuk mengklasifikasikan status gizi balita pada dataset penelitian. Agar perbedaan performa kedua algoritma dapat diamati secara lebih jelas, perbandingan nilai setiap metrik disajikan secara visual pada gambar berikut.



Gambar 4. Perbandingan Performa K-Nearest Neighbor dan Naïve Bayes

Sumber: Output Google Colab, 2026

KESIMPULAN DAN SARAN

Kesimpulan

Penerapan Cross-Industry Standard Process for Data Mining berhasil mengorganisasi proses klasifikasi status gizi balita mulai dari pemahaman kebutuhan, pemeriksaan karakteristik data, pembersihan dan transformasi, pemodelan, evaluasi, hingga rekomendasi implementasi. Dari 232 rekam awal diperoleh 218 data valid yang dibagi menjadi 174 data pelatihan dan 44 data pengujian. K-Nearest Neighbor dengan $K = 1$ menjadi model terbaik dengan akurasi 79,55%, presisi 57,81%, *recall* 59,19%, *F1-score* 58,25%, dan *area under the curve* 0,7418. Naïve Bayes menghasilkan akurasi 47,73%, presisi 31,56%, *recall* 30,06%, *F1-score* 29,41%, dan *area under the curve* 0,6914. Dengan demikian, K-Nearest Neighbor lebih adaptif untuk dataset lokal Posyandu Desa Pamanukan Kota dan dapat dijadikan dasar pengembangan sistem pendukung keputusan. Namun, kegagalan model mengenali kelas Gizi Lebih menunjukkan bahwa hasil klasifikasi tetap memerlukan verifikasi tenaga kesehatan.

Keterbatasan

Penelitian ini menggunakan dataset dari satu posyandu dan satu periode pengumpulan sehingga generalisasi hasil masih terbatas. Jumlah data pada setiap kelas tidak seimbang, terutama pada kelas Gizi Lebih, dan penelitian belum menerapkan teknik penyeimbangan seperti Synthetic Minority Over-sampling Technique. Evaluasi menggunakan satu skema pembagian *hold-out* 80:20 sehingga kestabilan model pada pembagian data lain belum diketahui. Variabel prediktor juga hanya mencakup jenis kelamin, umur, berat badan, dan tinggi badan, sedangkan faktor kesehatan, riwayat kelahiran, pola konsumsi, dan kondisi sosial ekonomi belum tersedia. Model yang dihasilkan masih berupa eksperimen pada Google Colaboratory dan belum diuji dalam aplikasi maupun alur kerja posyandu secara langsung.

Saran

Penelitian selanjutnya disarankan memperluas sumber data ke beberapa posyandu dan periode pengukuran, menambah jumlah data kelas minoritas, serta menguji Synthetic Minority Over-sampling Technique atau metode penyeimbangan lain. Penggunaan *cross-validation*, optimasi hiperparameter, dan perbandingan dengan algoritma lain dapat memberikan estimasi performa yang lebih stabil. Pengembangan sistem berbasis web atau perangkat seluler perlu dilengkapi validasi oleh ahli gizi, mekanisme pemeriksaan ulang, perlindungan data pribadi, dan uji kegunaan bersama kader. Dengan langkah tersebut, model dapat dikembangkan dari alat eksperimen menjadi sistem pendukung keputusan yang aman dan relevan bagi pelayanan kesehatan masyarakat.

Kontribusi Author

Penulis bertanggung jawab atas konseptualisasi penelitian, penyusunan kerangka teori, pengumpulan dan pembersihan data, pemodelan K-Nearest Neighbor dan Naïve Bayes, evaluasi hasil, interpretasi temuan, penulisan naskah, serta persetujuan terhadap versi akhir artikel.

REFERENSI

- Chapman, P. *et al.* (2000) *{CRISP-DM} 1.0: Step-by-Step Data Mining Guide*. Available at: <http://www.crisp-dm.org/CRISPWP-0800.pdf>.
- Dwiyanti, D.S. and Saptari, M.A. (2023) “Perancangan Sistem Informasi Status Gizi Balita (Stunting) di {UPTD} Puskesmas Muara Satu dan Muara Dua Menggunakan Metode {Naïve Bayes Classification} dan {K-Nearest Neighbor} Berbasis Web,” *Sisfo: Jurnal Ilmiah Sistem Informasi*, 7(2), p. 75. Available at: <https://doi.org/10.29103/sisfo.v7i2.14626>.
- Han, J., Kamber, M. and Pei, J. (2011) *Data Mining: Concepts and Techniques*. 3rd ed. Morgan Kaufmann. Available at: <https://doi.org/10.1016/C2009-0-61819-5>.
- Hardiani, T. and Putri, R.N. (2024) “Implementasi Metode {Naïve Bayes Classifier} untuk Klasifikasi Stunting pada Balita,” *Digital Transformation Technology*, 4(1), pp. 621–627. Available at: <https://doi.org/10.47709/digitech.v4i1.4481>.
- Indrisari, E., Febiansyah, H. and Adiwinoto, B. (2025) “A Systematic Literature Review on the Application of Machine Learning for Predicting Stunting Prevalence in Indonesia (2020--2024),” *Jurnal Sisfokom (Sistem Informasi dan Komputer)*, 14(3), pp. 277–283. Available at: <https://doi.org/10.32736/sisfokom.v14i3.2366>.
- Islam, H.I., Mulyadien, M.K. and Enri, U. (2022) “Penerapan Algoritma {C4.5} dalam Klasifikasi Status Gizi Balita,” *Jurnal Ilmiah Wahana Pendidikan*, 8(10), pp. 116–125. Available at: <https://doi.org/10.5281/zenodo.6791722>.
- Kementerian Kesehatan Republik Indonesia (2020) “Peraturan Menteri Kesehatan Republik Indonesia Nomor 2 Tahun 2020 tentang Standar Antropometri Anak.”
- Kementerian Kesehatan Republik Indonesia (2023) *Petunjuk Teknis Pemberian Makanan*

- Tambahan Berbahan Pangan Lokal bagi Ibu Hamil dan Balita*. Jakarta.
- Kotler, P., Kartajaya, H. and Setiawan, I. (2021) *Marketing 5.0: Technology for humanity*. Wiley.
- Kurniawan, H., Rahim, A. and Siswa, T. (2024) “Implementasi Algoritma {Gaussian Naïve Bayes} dalam Klasifikasi Status Gizi pada Balita,” *Building of Informatics, Technology and Science (BITS)*, 6(2), pp. 627–635. Available at: <https://doi.org/10.47065/bits.v6i2.5493>.
- Loka, S.K.P. and Marsal, A. (2023) “Comparison Algorithm of {K-Nearest Neighbor} and {Naïve Bayes Classifier} for Classifying Nutritional Status in Toddlers,” *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 3, pp. 8–14.
- Lonang, S. and Normawati, D. (2022) “Klasifikasi Status Stunting pada Balita Menggunakan {K-Nearest Neighbor} dengan Feature Selection {Backward Elimination},” *Jurnal Media Informatika Budidarma*, 6(1), pp. 49–56. Available at: <https://doi.org/10.30865/mib.v6i1.3312>.
- Nadhirah, S.A., Rosella, D. and Sari, K. (2024) “Uji validitas dan reliabilitas keseimbangan dinamis The Timed Up and Go Test pada pasien stroke,” *Jurnal Kesehatan Primer*, 9(2), pp. 121–130.
- Nurainun, N. *et al.* (2023) “Penerapan Algoritma {Naïve Bayes Classifier} dalam Klasifikasi Status Gizi Balita dengan Pengujian {K-Fold Cross Validation},” *Journal of Computer System and Informatics (JoSYC)*, 4(3), pp. 578–586. Available at: <https://doi.org/10.47065/josyc.v4i3.3414>.
- Pangumbara'an, M.S.G., Wakista, A.A. and Makhsun (2024) “Klasifikasi Status Gizi Balita Menggunakan Algoritma {K-Nearest Neighbor} dan {Naïve Bayes} (Studi Kasus: {UPTD} Puskesmas Bambu Apus),” *Infotech: Journal of Technology Information*, 10.
- World Health Organization (2025) *Global Report on Hypertension 2025: High Stakes: Turning Evidence into Action*. Available at: <https://www.who.int/publications/i/item/9789240115569>.
- World Health Organization (2026) “Tobacco and Nicotine.” Available at: <https://www.who.int/news-room/fact-sheets/detail/tobacco>.