



Perbandingan Kinerja Metode Naïve Bayes, SVM, dan ANN dalam Klasifikasi Status Gizi Balita

Chansa Adhista Syahira^{1*}, Arief Agoestanto^{2*}

¹ Mahasiswa Matematika, FMIPA, Universitas Negeri Semarang

² Matematika, FMIPA, Universitas Negeri Semarang

chansa.adhista08@gmail.com

Abstract

Nutritional problems in toddlers represent a crucial public health challenge. The increasing cases of undernutrition and overnutrition in toddlers necessitate data processing for classification to facilitate decision-making and preventive actions. Therefore, this study utilizes three methods, namely Naïve Bayes, SVM, and ANN, to evaluate their performance in classifying the nutritional status of toddlers. The analysis was conducted using the Google Colab platform through the stages of data collection, data preprocessing, data modeling, and performance evaluation using a confusion matrix, which serves to measure the performance of the machine learning classification models. The confusion matrix contains advanced evaluation metrics to assess model quality specifically, namely accuracy, precision, and recall. The results indicate that the ANN method, with an accuracy rate of 99%, is the most effective and superior method for monitoring toddler nutrition compared to SVM, which achieved 98% accuracy, and Naïve Bayes at 93%. This superiority is further supported by ANN's ability to iteratively adjust weights through hidden layers, enabling it to precisely model the non-linear relationships between the weight and height variables.

Keywords: classification; nutritional status of toddlers; Naïve Bayes; SVM; ANN

Abstrak

Masalah gizi pada balita merupakan tantangan kesehatan masyarakat yang krusial. Peningkatan kasus kekurangan dan kelebihan gizi pada balita memerlukan pengolahan data untuk proses klasifikasi sehingga dapat memberikan kemudahan dalam mengambil kesimpulan dan mengambil tindakan dalam pencegahan kekurangan dan kelebihan gizi balita. Oleh karena itu, penelitian ini menggunakan tiga metode yaitu Naïve Bayes, SVM, dan ANN untuk mengetahui kinerja dalam mengklasifikasikan status gizi balita menggunakan platform Google Colab dengan tahapan pengumpulan data, preprocessing data, pemodelan data, dan evaluasi performa berupa *confusion matrix* yang berfungsi untuk mengukur kinerja model klasifikasi machine learning. *confusion matrix* berisi metrik evaluasi lanjutan untuk menilai kualitas model secara spesifik yaitu *accuracy*, *precision*, dan *recall*. Hasil penelitian menunjukkan metode ANN dengan tingkat akurasi 99% merupakan metode terbaik dan paling efektif dalam pemantauan gizi balita dibandingkan SVM tingkat akurasi sebesar 98%, dan Naïve Bayes sebesar 93%. Hal ini juga didukung adanya keunggulan ANN dalam menyesuaikan bobot secara iteratif melalui hidden layers sehingga mampu memodelkan hubungan non-variabel antara variabel berat badan dan tinggi badan secara presisi.

Kata Kunci: klasifikasi; status gizi balita; Naïve Bayes; SVM; ANN

1. PENDAHULUAN

Di era modern ini perkembangan teknologi berupa komputer dan jaringan internet dijadikan sebagai sarana utama dalam penyampaian informasi juga dalam mendukung efisiensi pengolahan data di bidang pendidikan. Salah satu teknologi yang berperan krusial adalah data mining. *Data mining* mampu mengolah kumpulan data dalam jumlah besar menjadi informasi atau pengetahuan yang bermanfaat. Melalui penerapan *data mining*, pencarian pola maupun tren baru pada data berukuran besar dapat dilakukan dengan lebih mudah, sehingga dapat mendukung proses pengambilan keputusan, melakukan prediksi tertentu, serta memberikan analisis terkait langkah yang perlu diambil. Data mining juga dapat diterapkan pada klasifikasi data.

Klasifikasi merupakan operasi untuk memisahkan beragam entitas kedalam beberapa kategori (Mughnyanti & Hafiz Nanda Ginting, 2023). Pengklasifikasian merupakan pelatihan pada fungsi f (target/label) yang memasangkan setiap atribut x ke satu dari jumlah label kelas y yang tersedia. klasifikasi secara umum memberikan kemudahan dalam mengelompokkan data yang terstruktur juga pada data yang tidak terstruktur. Pada data yang tidak terstruktur, klasifikasi mempermudah pencarian dan pemanfaatan dengan memberikan label yang jelas. Selain itu, proses kategorisasi sangat berguna untuk membantu mengidentifikasi dan menghilangkan potensi adanya data yang terduplikasi secara tidak sengaja. Sehingga metode klasifikasi yang digunakan untuk analisis data adalah metode Naïve Bayes, SVM, dan ANN.

Naïve Bayes, SVM, dan ANN memiliki cara kerja yang berbeda. Naïve Bayes menggunakan teknik memprediksi yang berbasis probalistik sederhana (Ridwan, 2020), sementara SVM adalah model pembelajaran mesin yang mengoptimalkan fungsi linear dalam ruang fitur berdimensi tinggi berdasarkan prinsip teori pembelajaran statistik (Irawan et al., 2021), dan ANN adalah jaringan syaraf tiruan yang merupakan jaringan komputasi (system node dan interkoneksi antar node) (Sapanta, 2018).

Klasifikasi memiliki peran dalam bidang kesehatan termasuk kesehatan masyarakat karena dapat mengubah data mentah (seperti berat badan, tinggi badan, dan usia) menjadi informasi yang dapat ditindaklanjuti. Salah satu contoh dalam bidang kesehatan yaitu klasifikasi pada gizi balita. Klasifikasi status gizi balita merupakan langkah krusial dalam pengambilan keputusan medis dan kebijakan kesehatan. Implementasi ini memberikan manfaat strategis bagi berbagai pihak seperti Dinas Kesehatan dapat memetakan sebaran gizi secara real-time untuk intervensi wilayah, peneliti dapat mengidentifikasi faktor dominan penyebab masalah gizi, serta orang tua dapat memperoleh peringatan dini dan penyuluhan saat pertumbuhan anak mulai menjauhi kurva normal (Pertiwiningrum & Kamalah, 2021).

Peningkatan kasus gizi balita di Indonesia setiap tahun menjadikan para peneliti mulai melakukan penelitian untuk mendukung dalam upaya pencegahan stunting. Beberapa studi yang relevan dengan penelitian ini adalah penelitian menurut Hananti & Sari (2021) dengan judul "Perbandingan Metode *Support Vector Machine* (SVM), dan *Artificial*

Neural Network (ANN) pada Klasifikasi Gizi Balita”. Kesimpulan dari penelitian tersebut adalah metode ANN lebih besar dengan tingkat *accuracy* sebesar 94,82%, *precision* sebesar 51,00%, *recall* sebesar 51,09% sedangkan metode SVM tingkat *accuracy* sebesar 94,46%, *precision* sebesar 46,08%, *recall* sebesar 50,59%. Penelitian tersebut tidak menyediakan algoritma pembanding dasar (*baseline model*) yang berbasis probabilitas atau linear sederhana (seperti Naïve Bayes atau Regresi Logistik). Penelitian selanjutnya menurut Setiawan & Utama (2022) dengan judul "Klasifikasi Status Gizi Balita Menggunakan Metode Naïve Bayes Classifier". Pada penelitian ini menggunakan metode Naïve bayes classification dengan akurasi yang didapat sebesar 60%. Pengujian yang dilakukan pun menggunakan 5 Data sampel dengan metode pembagian data Fold Cross-Validation yang memberikan hasil rata-rata pada kategori BB/U adalah 80%, PB/U adalah 80%, dan BB/PB adalah 40%. Naïve Bayes mengasumsikan semua fitur bersifat independen (berdiri sendiri). Pada pertumbuhan manusia melibatkan korelasi yang sangat kuat antara umur, berat, dan tinggi badan. Hal ini menyebabkan Naïve Bayes kesulitan membaca pola pertumbuhan yang kompleks tersebut, sehingga akurasinya rendah.

Berdasarkan penelitian terdahulu, kebaruan penelitian ini menggunakan pendekatan evaluasi komparatif berskala luas secara head-to-head yang menggabungkan tiga algoritma yaitu Naïve Bayes, SVM, dan ANN. pendekatan evaluasi komparatif berskala luas secara head-to-head yang menggabungkan tiga algoritma secara umum hanya berfokus pada pengujian model tunggal, atau paling banyak melakukan komparasi biner yang memasang dua algoritma. Hal ini sering kali gagal memberikan gambaran yang utuh mengenai performa algoritma mana yang paling tangguh ketika dihadapkan pada satu karakteristik dataset yang sama. Dengan menggabungkan tiga algoritma penelitian ini mampu mengisi celah (*research gap*). Adanya evaluasi komparatif tiga arah ini membawa kebaruan berupa penyediaan standar metrik performa *accuracy*, *precision*, dan *recall* yang jauh lebih komprehensif, objektif, dan tepercaya untuk menentukan algoritma terbaik dalam menangani kompleksitas pola data gizi balita. Hasil prediksi pada tiga algoritma digunakan agar dapat melihat peluang naik atau turunnya status gizi untuk masa yang akan datang dan menjadikan hasil prediksi sebagai pedoman agar status gizi balita di Indonesia tetap dalam target yang ingin dicapai, serta memberikan masukan terhadap pihak terkait sehingga menjadi pertimbangan dalam menindaklanjuti pencegahan stunting (Putri et al., 2021).

2. METODE PENELITIAN

Pendekatan penelitian yang digunakan dalam penelitian ini adalah pendekatan kuantitatif. Kuantitatif adalah penelitian yang berfokus pada penggunaan angka sebagai ukuran datanya untuk menghasilkan kesimpulan yang dapat digeneralisasi. Jenis penelitian yang digunakan adalah penelitian kuantitatif komparatif. Penelitian komparatif bertujuan membandingkan keberadaan suatu variabel atau lebih pada dua atau lebih sampel yang berbeda, atau pada waktu yang berbeda. Metode komparatif

digunakan untuk mengetahui perbedaan keakuratan dari 3 metode yaitu *Naïve Bayes*, SVM, dan ANN pada sampel data gizi balita. Berikut langkah-langkah dalam melakukan analisis data:

1. Pengumpulan data

Penelitian ini menggunakan data gizi balita didapatkan melalui situs Kaggle yaitu (<https://www.kaggle.com/code/wansabrina/stunting-nutritional-data-eda-preprocessing>). Data berada di daerah Provinsi Sulawesi Selatan di Kabupaten Jeneponto, Kecamatan Bontoramba. Data yang digunakan pada tahun 2023 sebanyak 3808 data. Pada data mentah tersebut, dibutuhkan variabel status gizi balita sehingga data diolah terlebih dahulu dengan metode *clustering* yaitu algoritma K-Means untuk mendapatkan variabel status gizi balita berisi 5 kategori yaitu gizi buruk, gizi kurang, gizi baik, gizi lebih, dan obesitas.

Proses pengolahan K-Means menggunakan Microsoft Excel dengan melibatkan 3 fitur utama yaitu; umur, berat badan, dan tinggi badan. Proses awal K-Means dimulai dengan perhitungan normalisasi untuk variabel BB dan TB, kemudian menentukan titik pusat klaster awal (*Initial Centroid*) untuk kelima kelas tersebut. Proses iterasi dilakukan berulang dimana posisi centroid terus diperbarui dengan menghitung rata-rata (mean) baru dari data yang masuk ke klaster sehingga mencapai kondisi konvergen (titik pusat tidak berubah). Pada penelitian ini, iterasi dilakukan sebanyak 10 kali hingga centroid yang baru tidak mengalami perubahan pada centroid sebelumnya sehingga hasil akhir pada metode K-Means pada variabel status gizi balita didapatkan sebanyak 725 gizi baik, 1069 gizi buruk, 709 gizi kurang, 1106 gizi lebih, dan 199 obesitas.

Oleh karena itu variabel data yang akan digunakan dalam proses pengolahan klasifikasi yaitu; gender, umur, berat badan, tinggi badan, dan status gizi balita. Pada variabel gender dan status gizi balita termasuk ke dalam tipe data kategorikal sementara umur, berat badan, dan tinggi badan termasuk ke dalam tipe data numerik. Secara umum data gizi balita ini termasuk ke dalam data numerik kontinu. Hubungan variabel antara umur, BB, dan TB termasuk ke dalam prediktor utama sebagai antropometri anak, sementara gender sebagai prediktor pendukung atau pelengkap. Sementara status gizi ballita sebagai label yang akan kunci dalam mengklasifikasikan.

2. Preprocessing data

Pada tahapan ini data sudah terlebih dulu dilakukan pengecekan terhadap outlier atau kolom yang kosong sehingga tahapan dapat dilanjutkan dengan melakukan pembersihan dan persiapan data. Data gizi balita tidak memiliki meliputi:

a. Penentuan variabel

Dilakukan untuk mengetahui pembagian variabel x (input/prediktor) sebagai pengenalan pola hubungan antar data dan variabel y (output/target) sebagai label kategori dalam perbandingan dalam pengujian akurasi model.

b. Label kategori dalam bentuk numerik

Dilakukan agar data menjadi lebih terstruktur dan proses klasifikasi lebih efisien. Pada label encoding dimulai dari 0 karena dalam pemrograman menggunakan *zero-based indexing* artinya elemen pertama dalam sebuah array dihitung sebagai indeks 0. Hal ini memberikan efisien secara komputasi.

c. Pembagian data training dan data testing.

Dataset dibagi menjadi set pelatihan (*training set*) dan set pengujian (*testing set*) menggunakan rasio 80/20 untuk memastikan kemampuan generalisasi yang baik dan menghindari masalah *overfitting* (Martin et al., n.d.). Set pelatihan (80%) digunakan untuk mengajarkan pola-pola model di dalam data, sementara set pengujian (20%) mengevaluasi kinerja pada data yang belum pernah dilihat (*unseen data*). Distribusi *dataset* disajikan dalam Tabel 1

Tabel 1. Pembagian Dataset

Kategori	Jumlah Data	Persentase
Training	3046	80%
Testing	762	20%
Total	3808	100%

3. Pemodelan data

a. Normalisasi data

Normalisasi menggunakan metode standardization melalui fungsi *StandarScaler* untuk menyamakan skala antar fitur (Ramadhan & Khoirunnisa, 2021).

b. Pengujian model

Tahap ini, arsitektur model bermanfaat untuk membangun dan melatih data menggunakan data training sebesar 80% dari total data. Pada metode Naïve Bayes membangun model dilakukan dengan menginisialisasi fungsi *GaussianNB()* untuk menghitung probabilitas atau peluang awal dari setiap kategori status gizi balita. Pada pemodelan ini tidak ada parameter karena menggunakan nilai default bawaan dari *Scikit-Learn*, di mana parameter *priors* diatur ke *none* agar model secara mandiri dan objektif menghitung probabilitas awal berdasarkan proporsi kelas gizi balita yang riil dari data latihan. Selain itu, parameter *var_smoothing* pada bahasa pemrograman atau disebut sebagai nilai *default* sebesar 10^{-9} sudah bekerja dengan sangat baik dalam menjaga stabilitas perhitungan distribusi normal data antropometri balita. Sementara pada metode SVM dengan menentukan jenis kernel dengan RBF untuk menangani pola data yang tidak linear. Kernel SVM ini menggunakan 4 parameter, dan pada metode ANN menentukan jumlah lapisan (*hidden layers*) sebanyak 3 *hidden layers* dan jumlah neuron untuk memahami serta menangkap pola korelasi antara umur, berat, dan tinggi badan. Neuron dari *hidden layer* pertama adalah 32, sedangkan *hidden layer* kedua sebanyak 16, sementara pada *hidden layer* terakhir yaitu 5 dari total kelas status gizi balita. Epochs pada pemodelan ANN sebesar 100.

c. Prediksi

Tahapan selanjutnya adalah prediksi menggunakan data testing sebesar 20%. model akan menerima input berupa data balita baru yang telah dilakukan normalisasi dan mengeluarkan output berupa kategori status gizi balita. Hasil prediksi ini selanjutnya dilakukan perbandingan dengan data aktual untuk mengukur seberapa jauh model mampu bekerja di kondisi nyata.

d. Evaluasi dan Validasi

Untuk menilai kinerja dan akurasi metode Naïve Bayes, SVM, dan ANN dalam klasifikasi data gizi balita, tiga metrik evaluasi yang digunakan dalam penelitian ini. Pada klasifikasi *multiclass* dengan N kelas, *confusion matrix* akan berbentuk matriks persegi ukuran $N \times N$. Baris merepresentasikan kelas aktual (Sebenarnya), sedangkan kolom merepresentasikan kelas prediksi. Pada penelitian ini, klasifikasi *multiclass* untuk 5 kelas status gizi berbentuk seperti gambar berikut:

Aktual/Prediksi	Gizi Baik (C_1)	Gizi Buruk (C_2)	Gizi Kurang (C_3)	Gizi Lebih (C_4)	Obesitas (C_5)
Gizi Baik (C_1)	M_{11}	M_{12}	M_{13}	M_{14}	M_{15}
Gizi Buruk (C_2)	M_{21}	M_{22}	M_{23}	M_{24}	M_{25}
Gizi Kurang (C_3)	M_{31}	M_{32}	M_{33}	M_{34}	M_{35}
Gizi Lebih (C_4)	M_{41}	M_{42}	M_{43}	M_{44}	M_{45}
Obesitas (C_5)	M_{51}	M_{52}	M_{53}	M_{54}	M_{55}

Gambar 1. Klasifikasi *multiclass*

Pada gambar 2.1 diagonal utama ($M_{11}, M_{22}, M_{33}, M_{44}, M_{55}$) adalah jumlah data yang diprediksi dengan benar (aktual dan prediksi sama). Sementara nilai di luar diagonal adalah jumlah data yang salah prediksi (meleset ke kelas lain). Kemudian menghitung nilai TP, FP, TN, dan FN yang digunakan untuk mengetahui performa dari metode. Nilai-nilai ini memiliki makna sendiri yaitu nilai TP (*True Positive*) adalah jumlah data positif yang diprediksi benar, nilai FP (*False Positive*) adalah jumlah data positif yang diprediksi salah, nilai TN (*True Negative*) adalah jumlah data negatif yang diprediksi benar, dan nilai FN (*False Negative*) adalah jumlah data negatif yang diprediksi salah (Sathyanarayanan, 2024). Cara menghitung nilai-nilai tersebut adalah sebagai contoh pada kelas gizi baik (C_1). Nilai TP berada di diagonal utama maka $TP = M_{11}$. Nilai FP yaitu total dari kolom C_1 selain diagonal maka $FP = M_{21} + M_{31} + M_{41} + M_{51}$. Nilai FN yaitu total dari baris C_1 selain diagonal maka $FN = M_{12} + M_{13} + M_{14} + M_{15}$, dan nilai $TN = M_{22} + M_{33} + M_{44} + M_{55}$. Setelah nilai TP, FP, FN, dan TN didapatkan maka menghitung *accuracy*, *recall*, dan *precision*. *Accuracy* adalah menghitung keakuratan sistem mengklasifikasikan data dengan tepat (Ainurrohma, 2021). Rumus umum untuk *accuracy* adalah sebagai berikut pada eq. (1)

$$Accuracy = \frac{\sum TP}{\sum TP + TN + FP + FN} \quad (1)$$

Recall adalah ukuran kelengkapan kelesuruhan data yang bernilai positif (Azhari et al., 2021). Rumus umum untuk *recall* adalah sebagai berikut pada eq. (2)

$$Recall = \frac{TP}{TP+FN} \quad (2)$$

Precision dianggap sebagai ukuran ketepatan, contoh pada nilai persentase yang diberikan label positif (Cendani & Wibowo, 2022). Rumus umum untuk *precision* adalah sebagai berikut pada eq. (3)

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

Dari ketiga metrik evaluasi ini secara kolektif memberikan evaluasi komprehensif mengenai kinerja metode. *Accuracy* sebagai keakuratan sistem dalam mengklasifikasi secara tepat, *recall* sebagai ukuran kelengkapan keseluruhan data bernilai positif, sementara *precision* sebagai ukuran ketepatan.

3. HASIL DAN PEMBAHASAN

3.1 Hasil

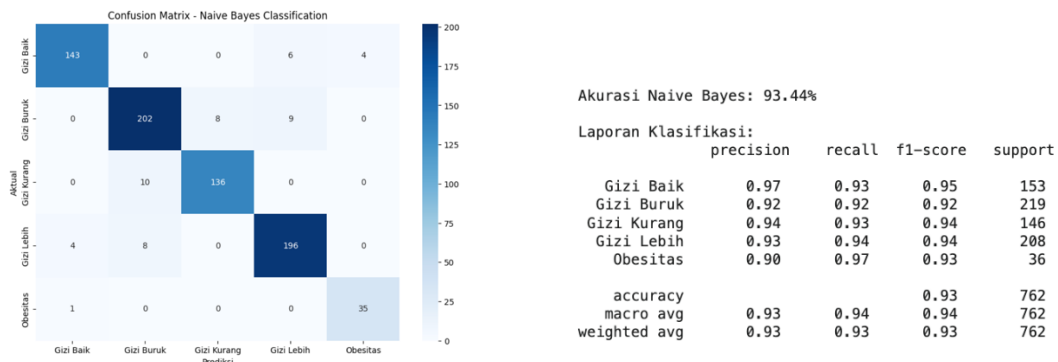
Pada data mentah yang telah diolah terlebih dahulu dengan metode *clustering* menggunakan algoritma K-Means, berikut adalah gambaran 10 data teratas.

No.	Gender	Age (Month)	Weight (BB)	Height (TB)	Status Gizi Balita
1	M	59	13	99	Gizi Kurang
2	M	28	11	83	Gizi Buruk
3	M	56	13,6	98	Gizi Kurang
4	F	59	12,6	97,2	Gizi Kurang
5	M	56	13,2	97	Gizi Kurang
6	F	54	12,6	96	Gizi Kurang
7	M	53	12,5	95	Gizi Kurang
8	F	56	14,5	98	Gizi Kurang
9	M	54	12,4	97	Gizi Kurang

Gambar 2. Data Gizi Balita

Data yang telah diolah dijadikan sebagai input bagi algoritma klasifikasi. Data ini mencakup parameter fisik utama yang menjadi acuan standar antropometri untuk menentukan kategori status gizi setiap balita. Data tersebut akan melalui proses normalisasi kemudian diolah dengan metode Naïve Bayes, SVM, dan ANN. Data diolah menggunakan Google Colab untuk memberikan kemudahan dalam perhitungan. Hasil pada masing-masing metode dengan *confusion matrix* dan klasifikasi sebagai berikut:

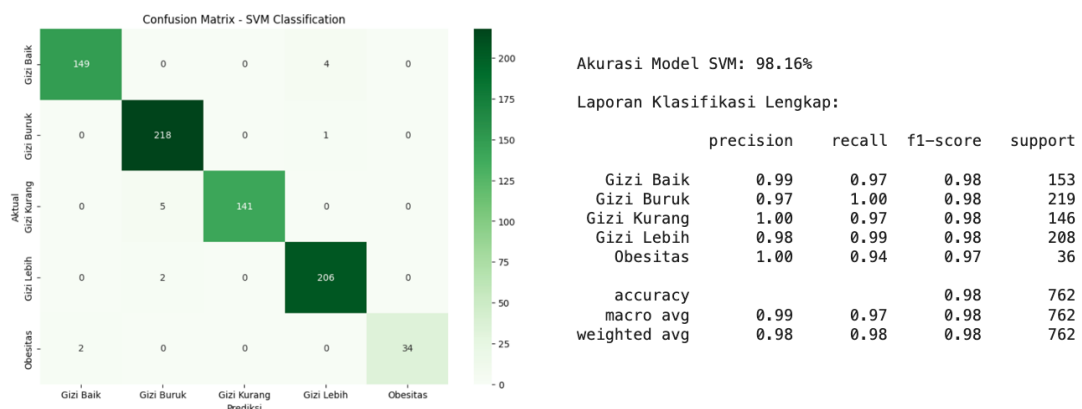
1. Naïve Bayes



Gambar 3. *Confusion Matrix* dan Hasil Klasifikasi Naïve Bayes

Analisis *confusion matrix* pada gambar 3.1 menunjukkan bahwa Naïve Bayes sangat akurat dalam mengklasifikasikan, karena hampir seluruh prediksi benar dengan nilai yang berada di luar diagonal *confusion matrix* sangat sedikit. Sementara pada gizi buruk menyatakan bahwa kinerja Naïve Bayes dalam membandingkan data aktual dengan data prediksi mengalami kesulitan (8 salah prediksi gizi kurang dan 9 prediksi sebagai gizi lebih). Hal ini mengindikasikan bahwa keterbatasan metode Naïve Bayes dalam membedakan batas-batas antar kategori memengaruhi akurasi perbandingan antara data aktual dan hasil prediksi. Sedangkan pada hasil klasifikasi menunjukkan tingkat akurasi pada metode Naïve Bayes sebesar 93%, sementara *precision* dengan nilai sebesar 93% dan *recall* sebesar 94%.

2. SVM

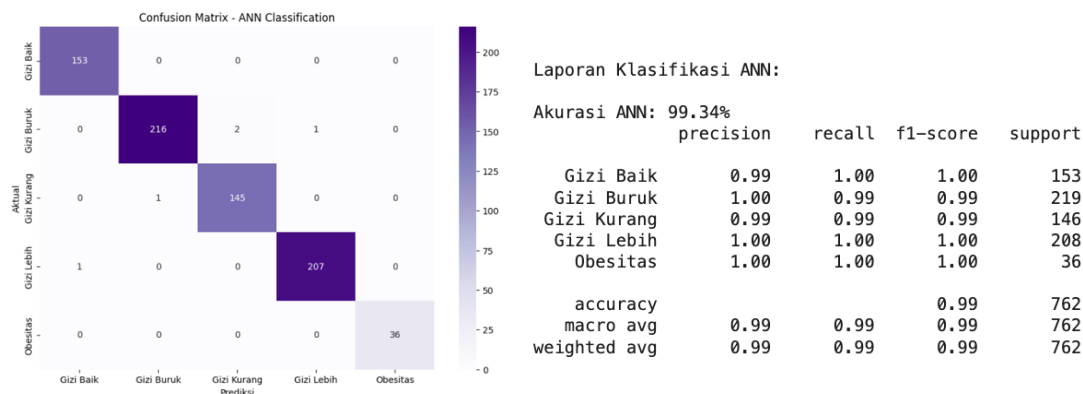


Gambar 4. *Confusion Matrix* dan Hasil Klasifikasi SVM

Analisis *confusion matrix* menunjukkan bahwa metode SVM sangat akurat dalam mengklasifikasikan kategori gizi. Meskipun terdapat beberapa kesalahan klasifikasi pada kategori gizi kurang dan gizi baik, secara keseluruhan model tetap menunjukkan kinerja yang kuat dengan mayoritas prediksi yang sesuai dengan data

aktual. Pada hasil klasifikasi menunjukkan tingkat akurasi pada metode SVM sebesar 98%, sementara *precision* sebesar 99% dan *recall* sebesar 98%.

3. ANN



Gambar 5. *Confusion Matrix* dan Hasil Klasifikasi ANN

Analisis *confusion matrix* menunjukkan bahwa metode ANN bekerja sangat optimal pada kategori gizi baik dan obesitas, dengan seluruh prediksi sesuai data aktual (kesalahan nol). Secara keseluruhan penerapan ANN menghasilkan akurasi prediksi tertinggi dibandingkan metode lainnya. Pada hasil klasifikasi menunjukkan tingkat akurasi pada metode ANN sebesar 99 %, sementara *precision* sebesar 99% dan *recall* sebesar 99%. Secara keseluruhan hasil pengujian tingkat *accuracy*, *recall*, dan *precision* pada kinerja metode Naïve Bayes, SVM, dan ANN dapat dilihat pada Tabel 2.

Tabel 2. Pengklasifikasian Data Gizi Balita

Metode	Accuracy	Precision	Recall
Naïve Bayes	93%	93%	94%
SVM	98%	99%	97%
ANN	99%	99%	99%

Berdasarkan hasil pengujian klasifikasi gizi balita, keakuratan dan keefektifan pada ANN lebih unggul dalam memahami karakteristik data yang kompleks dan bersifat non-linear dibandingkan SVM dan Naïve Bayes. Hal ini menunjukkan kemampuan ANN dalam memahami data yang kompleks sehingga pembaruan bobot (*weight adjustment*) selama proses pelatihan sehingga ANN lebih sensitif terhadap perbedaan kecil pada parameter fisik balita.

3.2 Pembahasan

Akurasi menunjukkan performa teknik klasifikasi secara keseluruhan. Semakin tinggi nilai akurasi maka semakin baik performa klasifikasi (Yudianto, 2020). Metode Naïve Bayes, SVM, dan ANN menunjukkan bahwa metode ANN memiliki tingkat akurasi yang sangat tinggi. Hal ini juga didukung dari hasil penelitian oleh (Hananti & Sari, 2021) yang membandingkan metode SVM dan ANN dengan hasil ANN adalah metode terbaik

dalam memprediksi kategori status gizi balita. Secara keseluruhan, evaluasi performa pada *confusion matrix* memperlihatkan bahwa metode ANN merupakan algoritma yang efektif dan efisien untuk klasifikasi status gizi balita. Hal ini juga didukung karena keunggulan ANN dalam memahami pola yang kompleks. Data gizi balita merupakan data yang bersifat multivariat atau banyak variabel dan memiliki pola non-linear yang kompleks sehingga metode dengan fleksibilitas tinggi seperti ANN mampu dalam memahami pola pertumbuhan. Dibandingkan dengan metode Naïve Bayes, metode ini kesulitan dalam memahami pola kompleks sedangkan SVM termasuk dalam kategori sedang dalam memahami pola non-linear. Dibalik keunggulan ANN dalam memahami pola kompleks, selama proses pengolahan data, ANN membutuhkan waktu cukup lama dibandingkan Naïve Bayes dan SVM. Oleh karena itu hasil penelitian ini menunjukkan bahwa ANN lebih cocok digunakan untuk analisis klasifikasi pada data serupa, dengan hasil tingkat akurasi dan kestabilan yang lebih tinggi.

4. SIMPULAN

Perbandingan performa ketiga algoritma terhadap accuracy, precision, dan recall pada ANN yaitu accuracy, precision, dan recall sebesar 99%. Sementara SVM yaitu accuracy sebesar 98%, precision 99%, dan recall 97%. Naïve Bayes dengan accuracy sebesar 93%, precision 93%, dan recall sebesar 94%. Berdasarkan metrik evaluasi yang diukur melalui platform Google Colab, Artificial Neural Network (ANN) terbukti menjadi metode yang paling efektif dalam klasifikasikan data gizi balita. Dengan demikian, penerapan metode ANN sangat direkomendasikan sebagai sistem pendukung keputusan bagi tenaga kesehatan di Kabupaten Jeneponto untuk mempercepat identifikasi masalah gizi secara efektif.

5. REKOMENDASI

Berdasarkan hasil penelitian yang telah dicapai, terdapat beberapa saran yang dapat diajukan untuk pengembangan penelitian di masa mendatang yaitu:

1. Disarankan menggunakan aplikasi data mining selain Google Colab seperti aplikasi RapidMiner untuk meminimalkan waktu selama proses pengolahan data.
2. Untuk penelitian selanjutnya, disarankan menambahkan fitur antropometri untuk meningkatkan generalisasi model.
3. Untuk penelitian selanjutnya, disarankan menambahkan penggunaan validasi silang (cross-validation)

6. REFERENSI

- Ainurrohma. (2021). Akurasi Algoritma Klasifikasi pada Software Rapidminer dan Weka. PRISMA, Prosiding Seminar Nasional Matematika, 4, 493–499. <https://journal.unnes.ac.id/sju/index.php/prisma/>
- Azhari, M., Situmorang, Z., & Rosnelly, R. (2021). Perbandingan Akurasi, Recall, dan Presisi Klasifikasi pada Algoritma C4.5, Random Forest, SVM dan Naive Bayes. Jurnal Media Informatika Budidarma, 5(2), 640. <https://doi.org/10.30865/mib.v5i2.2937>
- Cendani, L. M., & Wibowo, A. (2022). Perbandingan Metode Ensemble Learning pada Klasifikasi Penyakit Diabetes. Jurnal Masyarakat Informatika, 13(1), 33–44.

<https://doi.org/10.14710/jmasif.13.1.42912>

- Hananti, H., & Sari, K. (2021). Perbandingan Metode Support Vector Machine (SVM) dan Artificial Neural Network (ANN) pada Klasifikasi Gizi Balita Studi Kasus pada Puskesmas Salissingan (Comparison of Support Vector Machine (SVM) and Artificial Neural Network (ANN) Methods in Toddler Nut. Seminar Nasional Official Statistics, 2021(1), 1036–1043.
- Irawan, D., Perkasa, E. B., Yurindra, Y., Wahyuningsih, D., & Helmud, E. (2021). Perbandingan Klassifikasi SMS Berbasis Support Vector Machine, Naive Bayes Classifier, Random Forest dan Bagging Classifier. *Jurnal Sisfokom (Sistem Informasi Dan Komputer)*, 10(3), 432–437. <https://doi.org/10.32736/sisfokom.v10i3.1302>
- Martin, C. H., Peng, T. S., & Mahoney, M. W. (2021). Predicting trends in the quality of state-of-the-art neural networks without access to training or testing data. *Nature Communications*, 2021, 1–13. <https://doi.org/10.1038/s41467-021-24025-8>
- Mughnyanti, M., & Hafiz Nanda Ginting, S. (2023). Data Mining Manhattan Distance dan Euclidean Distance Pada Algoritma X-Means Dalam Klasifikasi Minat dan Bakat Siswa. *Remik*, 7(1), 835–842. <https://doi.org/10.33395/remik.v7i1.12162>
- Pertiwiningrum, D. A., & Kamalah, A. D. (2021). Prosiding Seminar Nasional Kesehatan 2021 Lembaga Penelitian dan Pengabdian Masyarakat Universitas Muhammadiyah Pekajangan Pekalongan. *Seminar Nasional Kesehatan*, 2017, 2148–2156.
- Putri, N. E., Andarini, M. Y., & Achmad, S. (2021). Gambaran Status Gizi pada Balita di Puskesmas Karang Harja Bekasi Tahun 2019. *Jurnal Riset Kedokteran*, 1(1), 14–18. <https://doi.org/10.29313/jrk.v1i1.108>
- Ramadhan, N. G., & Khoirunnisa, A. (2021). Klasifikasi Data Malaria Menggunakan Metode Support Vector Machine. *Jurnal Media Informatika Budidarma*, 5(4), 1580. <https://doi.org/10.30865/mib.v5i4.3347>
- Ridwan, A. (2020). Penerapan Algoritma Naïve Bayes Untuk Klasifikasi Penyakit Diabetes Mellitus. *Jurnal SISKOM-KB (Sistem Komputer Dan Kecerdasan Buatan)*, 4(1), 15–21. <https://doi.org/10.47970/siskom-kb.v4i1.169>
- Sapanta, M. D. (2018). Perbandingan Algoritma Pelatihan backpropagation pada Studi Peramalan Beban Menggunakan Metode Artificial neural network (ANN) Di Kabupaten Bantul. *Universitas Islam Indonesia*.
- Sathyanarayanan, S. (2024). Confusion Matrix-Based Performance Evaluation Metrics. *African Journal of Biomedical Research*, 27(4), 4023–4031. <https://doi.org/10.53555/ajbr.v27i4s.4345>
- Setiawan, H. B., & Utama, G. P. (2022). Klasifikasi Status Gizi Balita Menggunakan Metode Naïve Bayes Classifier. *Senafti*. 9(1), 707–715. <https://doi.org/10.31543/jii.v9i1.379>
- Yudianto, M. R. A. (2020). Wayang Dengan Algoritma Convolutional Neural Network. *Teknologi Informasi*, 4(2), 182–190. <https://www.academia.edu/download/93353033/386791048.pdf>