



Machine Learning Untuk Klasifikasi Gizi Balita Menggunakan Algoritma Random Forest

Taufik Hidayat^{1*}, Hanif Fajar Kurniawan², Asep Hardiyanto Nugroho³, Sukisno⁴, Robby Rizki⁵

^{1,2,3,4}*Teknik Informatika, Universitas Islam Syekh Yusuf,
Jl. Maulana Yusuf No.10, Tangerang, Banten 15118, Indonesia*

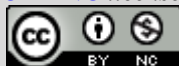
⁵*Teknik Informatika, Universitas Mathlaul Anwar,
Jl. Raya Labuan KM 23 Cikaliung, Pandeglang, Banten 42273, Indonesia*

*Email Penulis Koresponden: thidayat@unis.ac.id

Abstrak:

Kesehatan balita merupakan isu kritis dalam pembangunan suatu negara. Penilaian status gizi balita adalah langkah awal untuk mengidentifikasi risiko malnutrisi dan memberikan intervensi yang tepat. Dalam penelitian ini, kami mengusulkan sebuah pendekatan inovatif menggunakan teknik *Machine Learning*, khususnya algoritma *Random Forest*, untuk klasifikasi status gizi balita berdasarkan karakteristik demografis dan pola makan. *Dataset* yang digunakan terdiri dari informasi demografis seperti usia, jenis kelamin, berat badan, tinggi badan, dan data gizi pada setiap balita. Algoritma *Random Forest* dipilih karena kemampuannya dalam mengatasi *overfitting*, mengelola data yang tidak seimbang, dan memberikan hasil klasifikasi yang akurat. Berdasarkan penelitian yang telah dilakukan dapat ditarik kesimpulan Tingkat akurasi yang dihasilkan dari algoritma *random forest* sebesar 83% dari 168 sampel menunjukkan bahwa model klasifikasi yang digunakan memberikan prediksi yang sempurna atau benar untuk seluruh data uji yang digunakan.

This is an open access article under the [CC BY-NC](#) license



Kata Kunci:

*Stunting;
Random Forest;
Gizi Balita;
Klasifikasi;*

Riwayat Artikel:

Diserahkan 11 Juli 2023
Direvisi 31 Juli 2024
Diterima 03 Juni 2025

DOI:

10.22441/incomtech.v15i2.30517

1. PENDAHULUAN

Stunting merupakan kondisi kerdil pada balita saat lahir, terus menjadi isu utama dalam kesehatan gizi di Indonesia. Menurut data yang dikumpulkan oleh *World Health Organization (WHO)*, Indonesia termasuk salah satu dari tiga negara dengan prevalensi stunting tertinggi di kawasan Asia Tenggara. Antara tahun 2005 hingga 2017, rata-rata prevalensi stunting pada balita di Indonesia mencapai 36,4%. Hal ini menjadi masalah gizi yang signifikan di Indonesia, terbukti dari hasil Pemantauan Status Gizi (PSG) yang menunjukkan Stunting pada anak usia dini memiliki tingkat kejadian yang lebih tinggi dibandingkan masalah gizi lainnya seperti kekurangan

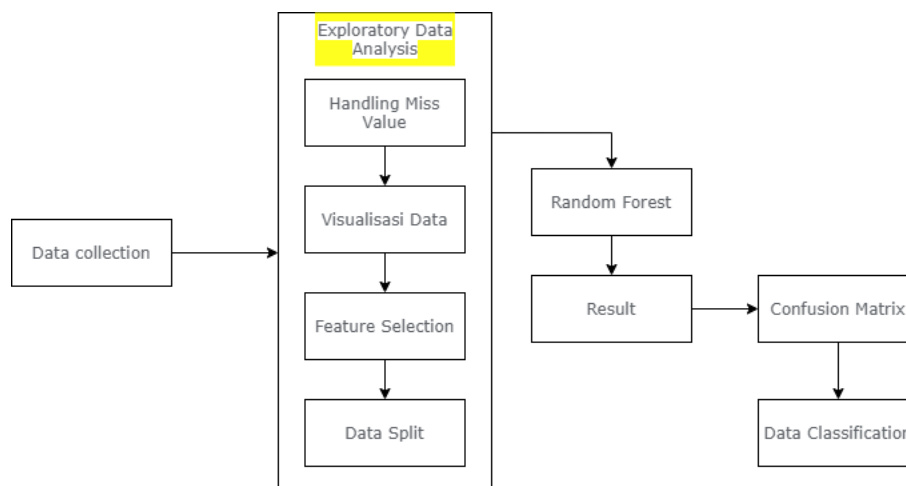
gizi, kurus, dan obesitas. Tingkat stunting pada anak usia dini juga meningkat dari 27,5% pada tahun 2016 dan menjadi 29,6% pada tahun 2017 [1], [2], [3], [4].

Kurangnya pertumbuhan pada anak kecil adalah masalah gizi yang serius dalam bidang kesehatan global. Anak-anak yang mengalami stunting dapat mengalami penurunan kemampuan berpikir, kesulitan berbicara, dan kesulitan dalam menyerap materi pelajaran secara efektif. Penyebab stunting pada anak bisa terjadi selama masa kehamilan, kelahiran, masa menyusui, atau fase setelah kelahiran. Ketidakcukupan nutrisi dari makanan pendamping ASI (MPASI) menjadi salah satu faktor utama. Selain itu, lingkungan yang tidak bersih juga bisa menyebabkan anak mudah terserang penyakit, dan pola asuh yang kurang baik juga berperan sebagai penyebab stunting [5], [6], [7].

Penting untuk diingat bahwa melihat kondisi fisik saja tidaklah cukup untuk menilai status gizi anak. Keterlibatan orang tua, terutama ibu, dalam memberikan makanan yang tepat sangatlah penting pada fase ini. Pola makan yang diberikan memiliki dampak besar pada pertumbuhan dan perkembangan anak, mengingat pentingnya asupan gizi yang mencukupi pada periode ini [8], [9], [10].

Penelitian ini menggunakan algoritma *random forest* untuk memprediksi keadaan stunting pada balita dengan melakukan pengolahan data terhadap 837 record data yang terdiri dari 11 atribut dan 1 kelas atribut, sehingga memperoleh hasil terbaik dalam memprediksi menggunakan algoritma tersebut.

2. METODE



Gambar 1. Flowchart usulan penelitian

Metode merupakan sebuah proses yang dilakukan untuk memecahkan sebuah masalah yang diangkat dalam sebuah penelitian sehingga akan ditemukan hasil yang akurat dan hingga dapat diambil kesimpulan. *Dataset* yang digunakan adalah *dataset* stunting dengan jumlah sebanyak 837 *record* data yang terdiri dari 11 atribut dan 1 kelas atribut, dengan model yang digunakan adalah Algoritma *Random Forest* [11], [12], [13], [14], [15], [16].

Random forest dimulai dengan menggunakan teknik dasar data mining, yaitu *decision tree*. Proses ini melibatkan penggunaan *decision tree* di mana input diterapkan dari bagian atas (*root*) dan diolah ke bagian bawah (*leaf*) untuk menentukan kelas pada data tersebut. Secara umum, *random forest* merupakan

serangkaian pengklasifikasi *decision tree* yang bekerja bersama-sama, di mana setiap *tree* memberikan suara terhadap kelas yang paling umum untuk input x . Dengan demikian, *random forest* dapat dianggap sebagai kumpulan *decision tree* yang berkolaborasi untuk mengklasifikasikan data ke dalam kelas-kelas tertentu [17], [18], [19], [20].

3. HASIL DAN PEMBAHASAN

3.1 Data Analisis / EDA.

Total *dataset* dalam penelitian ini meliputi 837 data *record* berupa informasi stunting pada balita. Sedangkan variabel input penelitian adalah 1) *ISO code* (Negara Asal), 2) Jenis Kelamin, 3) Usia, 4) Tinggi Badan, 5) Berat Badan, 6) Benua, 7) Kalangan, 8) Gizi Buruk, 9) Kelebihan Gizi, 10) Stunting, dan 11) Kekurangan Gizi dan mempunyai 1 kelas atribut yaitu status. Dapat dilihat Tabel 1 menunjukkan sampel *dataset* untuk digunakan dalam pengujian algoritma.

Tabel 1. Sempel *Dataset*

No	ISO Code (Negara Asal)	JK/Jenis Kelamin	Usia (Tahun)	Tinggi Badan (CM)		
0	AFG	1	3	83		
1	AFG	0	3	84		
2	AFG	0	5	75		
3	ALB	0	5	96		
4	ALB	0	3	75		

Berat Badan (KG)	Benua	Kalangan	Gizi Buruk	Kelebihan Gizi	Stunting
13	Asia	Miskin	18.2	6.5	53.2
14	Asia	Miskin	8.6	4.6	59.3
16	Asia	Miskin	9.5	5.4	40.9
12	Eropa	Menengah Keatas	8.1	9.5	20.4
18	Eropa	Menengah Keatas	12.2	30.0	39.2

Kekurangan Gizi	Status
44.9	Stunting
32.9	Stunting
25.0	Stunting
20.4	Stunting
39.2	Stunting

3.1 Testing Missing Value dan Pembersihan Missing Value.

Selanjutnya adalah proses *testing missing value*. Ditemukan sekitar 246 data *missing value* pada 4 atribut yaitu, 1) Gizi Buruk sebanyak 44 data, 2) Kelebihan Gizi sebanyak 513 data, 3) Stunting sebanyak 34 data, dan 4) Kekurangan Gizi sebanyak 15 data. Dapat dilihat Tabel 2 menunjukkan *missing value* pada *dataset*.

Tabel 2. *Dataset Missing Value*

Sensitivitas	Waktu Responden
ISO code (Negara Asal)	0
JK	0
Usia (Tahun)	0
TB (cm)	0
BB (kg)	0
Benua	0
Kalangan	0
Gizi_Buruk	44
Kelebihan_Gizi	153
Stunting	34
Kekurangan_Gizi	15
Status	0
Dtype: int64	

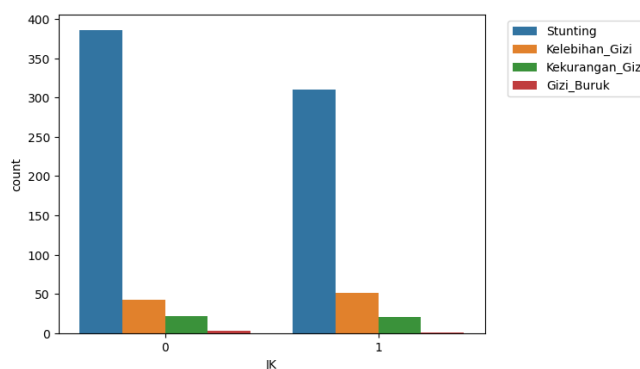
Ditemukan *missing value* sebesar 246 data pada *dataset*. Selanjutnya yaitu pembersihan pada *dataset* sehingga menghasilkan *dataset* yang bersih. Dapat dilihat Tabel 3 menunjukkan *dataset* setelah pembersihan *missing value*.

Tabel 3. Hasil Pembersihan

Sensitivitas	Waktu Responden
ISO code (Negara Asal)	0
JK	0
Usia (Tahun)	0
TB (cm)	0
BB (kg)	0
Benua	0
Kalangan	0
Gizi_Buruk	0
Kelebihan_Gizi	0
Stunting	0
Kekurangan_Gizi	0
Status	0
Dtype: int64	

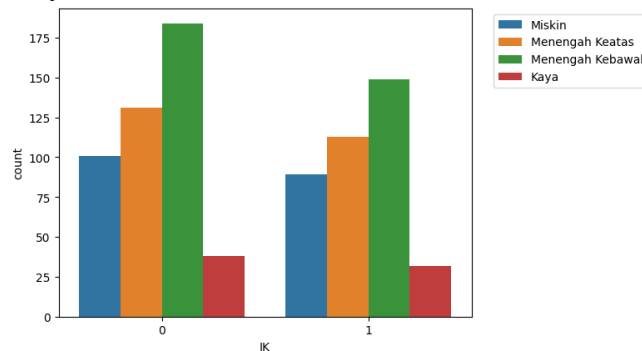
3.2 Count Plot.

Sebuah *count plot* yang menampilkan jumlah pengamatan untuk setiap jenis kelamin dengan pemisahan berdasarkan status. Dapat dilihat pada Gambar 1 yang menjelaskan *count plot* berdasarkan jenis penyakit.

Gambar 1. *Count Plot* Berdasarkan Jenis Penyakit

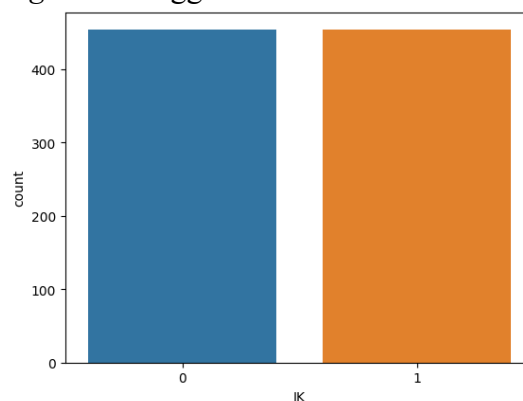
Berdasarkan gambar 1 tersebut dapat dijelaskan bahwa angka stunting menunjukkan jumlah tertinggi dibandingkan dengan kategori lainnya. Perhitungan ini diukur berdasarkan jenis kelamin (0 = Perempuan, 1 = Laki - Laki) yang dari keduanya menunjukkan angka stunting yang cukup tinggi.

Dapat dilihat pada Gambar 2 berikut ini menjelaskan perhitungan ini diukur berdasarkan jenis kelamin yang menunjukkan bahwa jumlah angka kalangan keluarga terbanyak berada di kalangan menengah kebawah diikuti oleh kalangan menengah keatas, selanjutnya kalangan miskin dan yang berada di urutan terakhir adalah kalangan kaya.



Gambar 2. *Count Plot* Berdasarkan Kalangan

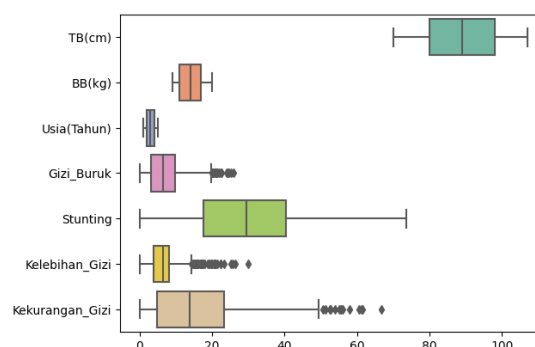
Dan dapat dilihat pada Gambar 3 berikut ini yang dilakukan yaitu *balancing* data, dimana sebelumnya data masih terlihat tidak beraturan, jadi agar terlihat setara dilakukanlah *balancing* data menggunakan metode *smote*.



Gambar 3. *Count Plot* Hasil Jumlah Pengamatan Berdasarkan Kategori

3.3 Box Plot.

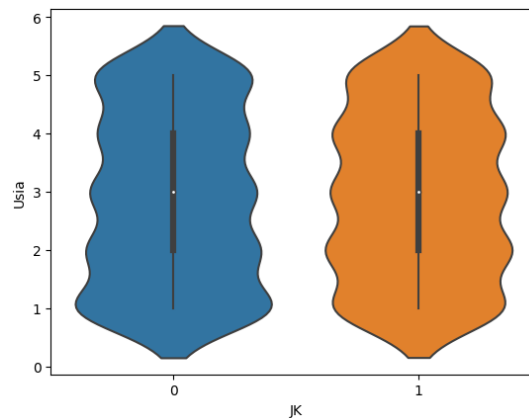
Box Plot digunakan untuk membuat subset data dengan fitur yang tercantum dalam variabel 'data1', menggunakan *Seaborn* untuk *visualisasi*. Dapat dilihat pada gambar 4.



Gambar 4. Box Plot

3.4 Violin Plot.

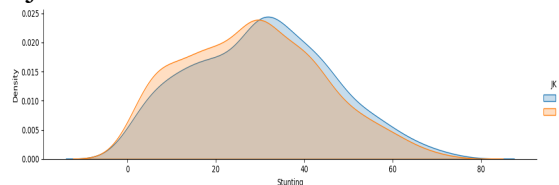
Violin Plot akan menunjukkan distribusi data berdasarkan jenis kelamin (JK) terhadap status stunting, seperti pada gambar dibawah ini , perempuan (0) yang terkena stunting dominan di usia 1 tahun, dan laki-laki (1) pada usia 2 tahun. Dapat dilihat pada Gambar 5 menunjukkan *violin plot* berdasarkan jenis kelamin dan usia.



Gambar 5. Violin Plot Berdasarkan Jenis Kelamin dan Usia

3.5 FacetGrid.

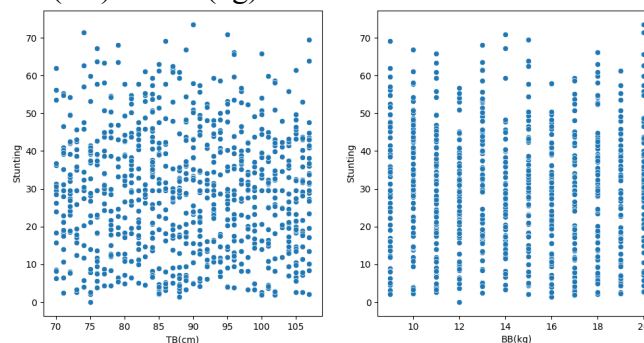
Memvisualisasikan distribusi data “Stunting” berdasarkan kategori "JK" secara terpisah dalam satu grid plot dengan membedakan warna dan menyertakan legenda untuk identifikasi yang lebih baik. Dapat dilihat pada Gambar 6 yang menjelaskan *facetgrid* berdasarkan jenis kelamin.



Gambar 6. FacetGrid Berdasarkan Jenis Kelamin

3.6 Scatter Plot.

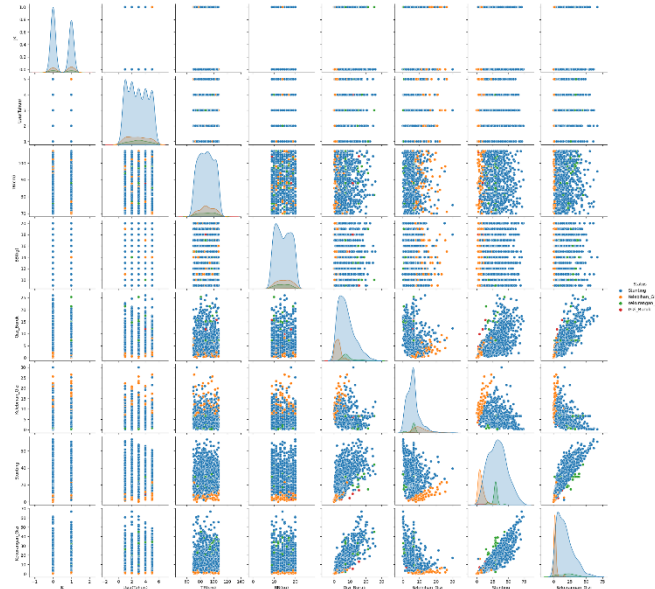
Memvisualisasikan hubungan antara tinggi badan ('TB') atau berat badan ('BB') terhadap stunting, yang mungkin memungkinkan untuk melihat pola atau korelasi antara variabel tersebut. Sebagai contoh Tinggi Badan 105cm dominan yang terkena stunting pada 20. Dapat dilihat pada Gambar 7 yang menjelaskan *scatter plot* berdasarkan tb(cm) dan bb(kg).



Gambar 7. Scatter Plot Berdasarkan TB(cm) dan BB(KG)

3.7 Pair Plot.

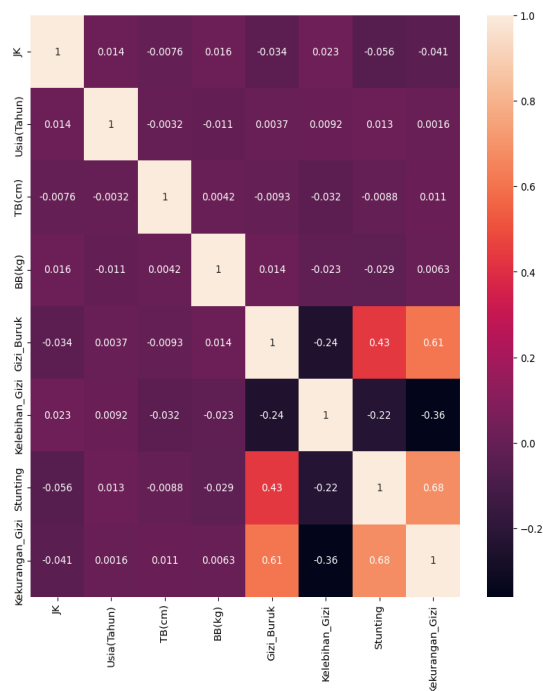
Membuat *grid* dari *scatter plot* antara pasangan variabel dalam *dataset*, dengan warna yang berbeda-beda sesuai dengan kategori yang diberikan dalam kolom 'JK'. Dapat dilihat pada Gambar 8 berikut ini.



Gambar 8. Pair Plot

3.8 Plot Figure.

Heatmap yang menunjukkan seberapa kuat korelasi antara pasangan variabel dalam *dataset*, dimana warna dan angka di dalam kotak-kotak *heatmap* akan memberikan informasi tentang tingkat korelasi antara variabel-variabel tersebut. Dapat dilihat pada gambar 9 sebagai berikut.



Gambar 9. Plot Figure

Dari gambar 9 dapat dijelaskan hubungan antara satu variabel dengan variabel lainnya, seperti misalnya, $x=JK$ dan $y=JK$ dari hasil hubungan keduanya memiliki nilai 1 dikarenakan sesuai. Tetapi jika $x=JK$ dan $y=usia$ hasilnya adalah 0,014, dikarenakan jarak korelasi antar keduanya berbeda.

3.9 Menghapus kolom / *feature* yang tidak penting.

Fungsi dari menghapus kolom – kolom tersebut adalah untuk membersihkan *dataset* yang dianggap tidak relevan atau tidak dibutuhkan untuk analisis atau pemodelan data yang sedang dilakukan. Dapat dilihat pada Tabel 4 berikut ini.

Tabel 4. Hasil Penghapusan Kolom / *Feature*

No	JK/ Jenis Kelamin	Usia (Tahun)	Tinggi Badan (CM)	Berat Badan (KG)
0	1	3	83	13
1	0	3	84	14
2	0	5	75	16
3	0	5	96	12
4	0	3	75	18
...	...	...	...	...
832	0	4	73	20
833	0	5	100	10
834	0	2	90	20
835	1	1	87	20
836	1	3	81	12

Gizi Buruk	Kelahiran Gizi	Stunting	Kekurangan Gizi	Status
18.2	6.5	53.2	44.9	Stunting
8.6	4.6	59.3	32.9	Stunting
9.5	5.4	40.9	25.0	Stunting
8.1	9.5	20.4	7.1	Stunting
12.2	30.0	39.2	17.0	Stunting
...	...	...	...	...
7.3	9.1	35.8	14.0	Stunting
3.8	3.5	35.1	12.7	Stunting
3.1	5.8	32.3	10.1	Stunting
3.3	3.6	27.6	11.2	Stunting
3.2	5.6	26.8	8.4	Stunting

3.10 *Split Training* dan *Testing Data*.

Untuk membagi data menjadi data latih dan data uji kami membuat 2 kode dibawah ini. Dan telah melatih model klasifikasi *Random Forest* menggunakan data latih tersebut. Model yang telah dilatih ini siap digunakan untuk melakukan prediksi pada data uji (X_{test}).

```
from sklearn.model_selection import train_test_split
X_train, X_test, y_train, y_test = train_test_split(predictors, target, test_size = 0.2,
random_state=42)
```

Rumus *Random Forest* yaitu sebagai berikut :

$$MSE = \frac{1}{N} \sum_{i=1}^N (f_i - y_i)^2$$

MSE = Mean Squared Error, merupakan suatu metrik pengukuran kesalahan prediksi model.

N = Jumlah total data dalam set.

f_i = Prediksi atau nilai yang dihasilkan oleh model untuk data ke-*i*

y_i = Nilai sebenarnya dari data ke-*i*

$$GINI = 1 - \sum_{i=1}^C (P_i)^2$$

GINI = indeks Gini

C = jumlah kelas atau kelompok dalam distribusi

P_i = proporsi kumulatif dari jumlah penduduk atau penghasilan dari kelompok ke-*i* dalam distribusi.

$$Entropy = \sum_{i=1}^C - P_i \log_2(P_i)$$

Entropy = nilai entropi dari sistem atau data.

P_i = probabilitas masing – masing kejadian ke-*i*

C = kejadian yang mungkin dalam sistem

$\log_2 (P_i)$ = menghitung logaritma basis 2 dari *P_i*

$- P_i \log_2 (P_i)$ = merupakan kontribusi entropi dari setiap kejadian ke-*i*.

3.11 Prediksi Model.

Fungsi dari model.predict(x_test) adalah bagian dari proses pemodelan untuk melakukan prediksi menggunakan model yang telah dilatih sebelumnya. Dapat dilihat pada Tabel 5.

Tabel 5. Prediksi Model

‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘kelebihan gizi’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘kelebihan gizi’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Kelebihan gizi’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘kelebihan gizi’,
‘Stunting’,	‘Stunting’,	‘Kelebihan gizi’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘kelebihan gizi’,
‘Stunting’,	‘Stunting’,	‘Kelebihan gizi’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘kelebihan gizi’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘kelebihan gizi’,	‘kekurangan gizi’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘kelebihan gizi’,
‘kelebihan gizi’,	‘Stunting’,	‘Stunting’,	‘Stunting’,	‘Stunting’,
‘Stunting’,	‘Stunting’,	‘Stunting’,	‘kekurangan gizi’,	‘Stunting’,

'Stunting',	'Stunting',	'Stunting',	'Stunting',	'kekurangan gizi',
'Stunting',	'kelebihan gizi',	'kekurangan gizi',	'Stunting',	'Stunting',
'kelebihan gizi',	'Stunting',	'Stunting',	'Stunting',	'Stunting',

3.12 Hasil dari Akurasi Model.

Hasil dari kedua fungsi ini memberikan *insight* penting tentang performa model klasifikasi dalam mengidentifikasi kelas-kelas pada data uji. Matriks konfusi memberikan gambaran tentang seberapa baik model mengklasifikasikan setiap kelas secara spesifik, sementara *classification report* memberikan ringkasan statistik yang mencakup *precision*, *recall*, dan *f1-score* untuk setiap kelas. Di samping itu, *accuracy_score* memberikan nilai *persentase* prediksi yang benar secara keseluruhan. didapatkan akurasi sebesar 83% dari 3 status yaitu kekurangan gizi, kelebihan gizi dan *Stunting* dengan total 168 sampel. Dapat dilihat hasil dari akurasi model pada Tabel 6.

Tabel 6. Akurasi Model

	precision	recall	f1-score	support
0	0.33	0.16	0.22	25
1	0.87	0.94	0.90	143
accuracy			0.83	168
Macro avg	0.60	0.55	0.56	168
Weighted avg	0.79	0.83	0.80	168

4. KESIMPULAN

Berdasarkan penelitian yang telah dilakukan dengan judul Penerapan Klasifikasi Stunting Pada Balita Menggunakan Algoritma *Random Forest*: Studi Kasus Pada Data Kesehatan Anak Usia Dini, dapat ditarik kesimpulan Tingkat akurasi yang dihasilkan dari algoritma *random forest* sebesar 83% dari 168 sampel menunjukkan bahwa model klasifikasi yang digunakan memberikan prediksi yang sempurna atau benar untuk seluruh data uji yang digunakan.

Penting untuk dievaluasi lebih dalam terhadap model yang digunakan untuk memastikan hasil yang diperoleh bukan karena *overfitting*, bisa melakukan pengujian *cross-validation*, memeriksa performa model pada *dataset* yang mencakup lebih luas.

REFERENSI

- [1] I. Choliq, D. Nasrullah, and M. Mundakir, "Pencegahan Stunting di Medokan Semampir Surabaya Melalui Modifikasi Makanan Pada Anak," *Humanism : Jurnal Pengabdian Masyarakat*, vol. 1, no. 1, Apr. 2020, doi: 10.30651/hm.v1i1.4544.
- [2] A. Sandi, K. Kusriani, and K. Kusnawi, "Analisa Prediksi Kesejahteraan Masyarakat Nelayan Lombok Timur Menggunakan Algoritma Random Forest," *Infotek : Jurnal Informatika dan Teknologi*, vol. 6, no. 2, pp. 238–248, Jul. 2023, doi: 10.29408/jit.v6i2.10104.

- [3] Fitriani *et al.*, “Cegah Stunting Itu Penting!,” *Jurnal Pengabdian Kepada Masyarakat (JurDikMas) Sosiosaintifik*, vol. 4, no. 2, pp. 63–67, Aug. 2022, doi: 10.54339/jurdikmas.v4i2.417.
- [4] R. D. Widjayatri, Y. Fitriani, and B. Tristyanto, “Sosialisasi Pengaruh Stunting Terhadap Pertumbuhan dan Perkembangan Anak Usia Dini,” *Murhum : Jurnal Pendidikan Anak Usia Dini*, pp. 16–27, Nov. 2020, doi: 10.37985/murhum.v1i2.11.
- [5] S. Anwar, E. Winarti, and S. Sunardi, “SYSTEMATIC REVIEW FAKTOR RISIKO, PENYEBAB DAN DAMPAK STUNTING PADA ANAK,” *Jurnal Ilmu Kesehatan*, vol. 11, no. 1, p. 88, Dec. 2022, doi: 10.32831/jik.v11i1.445.
- [6] R. D. P. Utami, “POLA PEMBERIAN MAKAN, PEMBERIAN ASI EKSklusif, ASUPAN PROTEIN DAN ENERGI, SEBAGAI PENYEBAB STUNTING DI DESA GROGOL PONOROGO,” *Jurnal Keperawatan Malang*, vol. 5, no. 2, pp. 96–102, Dec. 2020, doi: 10.36916/jkm.v5i2.114.
- [7] M. D. Khairani, K. Tjahjono, A. Rosidi, A. Margawati, and E. R. Noer, “Faktor determinan riwayat kehamilan dan kelahiran sebagai penyebab stunting,” *AcTion: Aceh Nutrition Journal*, vol. 8, no. 1, p. 70, Mar. 2023, doi: 10.30867/action.v8i1.793.
- [8] E. Fitriahadi *et al.*, “Meningkatkan Pengetahuan dan Kesadaran Tentang Stunting Sebagai Upaya Pencegahan Terjadinya Stunting,” *Jurnal Masyarakat Madani Indonesia*, vol. 2, no. 4, pp. 410–415, Oct. 2023, doi: 10.59025/js.v2i4.154.
- [9] C. Christine, F. V. M. Politon, and F. Hafid, “Sanitasi rumah dan stunting di Wilayah Kerja Puskesmas Labuan Kabupaten Donggala,” *AcTion: Aceh Nutrition Journal*, vol. 7, no. 2, p. 146, Nov. 2022, doi: 10.30867/action.v7i2.536.
- [10] W. Angraini, F. Firdaus, B. A. Pratiwi, O. Oktarianita, And H. Febriawati, “Pola Asuh, Pola Makan Dan Kondisi Lingkungan Fisik Dengan Kejadian Stunting,” *Journal of Nursing and Public Health*, vol. 11, no. 2, pp. 500–511, Oct. 2023, doi: 10.37676/jnph.v11i2.5186.
- [11] A. Ramadhan, B. Susetyo, and Indahwati, “Penerapan Metode Klasifikasi Random Forest Dalam Mengidentifikasi Faktor Penting Penilaian Mutu Pendidikan,” *Jurnal Pendidikan dan Kebudayaan*, vol. 4, no. 2, pp. 169–182, Dec. 2019, doi: 10.24832/jpnk.v4i2.1327.
- [12] M. S. Haris, M. Anshori, and A. N. Khudori, “PREDICTION OF STUNTING PREVALENCE IN EAST JAVA PROVINCE WITH RANDOM FOREST ALGORITHM,” *Jurnal Teknik Informatika (Jutif)*, vol. 4, no. 1, pp. 11–13, Feb. 2023, doi: 10.52436/1.jutif.2023.4.1.614.
- [13] R. Setiawan and A. Triayudi, “Klasifikasi Status Gizi Balita Menggunakan Naïve Bayes dan K-Nearest Neighbor Berbasis Web,” *Jurnal Media Informatika Budidarma*, vol. 6, no. 2, p. 777, Apr. 2022, doi: 10.30865/mib.v6i2.3566.
- [14] R. Aryanti, T. Misriati, and A. Sagiyanto, “Analisis Sentimen Aplikasi Primaku Menggunakan Algoritma Random Forest dan SMOTE untuk Mengatasi Ketidakseimbangan Data,” *Journal of Computer System and Informatics (JoSYC)*, vol. 5, no. 1, pp. 218–227, Nov. 2023, doi: 10.47065/josyc.v5i1.4562.
- [15] I. P. Putri, T. Terttiaavini, and N. Arminarahmah, “Analisis Perbandingan Algoritma Machine Learning untuk Prediksi Stunting pada Anak,” *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 1, pp. 257–265, Jan. 2024, doi: 10.57152/malcom.v4i1.1078.
- [16] H. Saleh, M. Faisal, and R. I. Musa, “KLASIFIKASI STATUS GIZI BALITA MENGGUNAKAN METODE K-NEAREST NEIGHBOR,” *Simtek : jurnal sistem informasi dan teknik komputer*, vol. 4, no. 2, pp. 120–126, Oct. 2019, doi: 10.51876/simtek.v4i2.60.
- [17] Oon Wira Yuda, Darmawan Tuti, Lim Sheih Yee, and Susanti, “Penerapan Penerapan Data Mining Untuk Klasifikasi Kelulusan Mahasiswa Tepat Waktu Menggunakan Metode Random Forest,” *SATIN - Sains dan Teknologi Informasi*, vol. 8, no. 2, pp. 122–131, Dec. 2022, doi: 10.33372/stn.v8i2.885.
- [18] Yoga Religia, Agung Nugroho, and Wahyu Hadikristanto, “Klasifikasi Analisis Perbandingan Algoritma Optimasi pada Random Forest untuk Klasifikasi Data Bank Marketing,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 1, pp. 187–192, Feb. 2021, doi: 10.29207/resti.v5i1.2813.

- [19] G. A. Sandag, "Prediksi Rating Aplikasi App Store Menggunakan Algoritma Random Forest," *CogITO Smart Journal*, vol. 6, no. 2, pp. 167–178, Dec. 2020, doi: 10.31154/cogito.v6i2.270.167-178.
- [20] N. Widjiyati, "Implementasi Algoritme Random Forest Pada Klasifikasi Dataset Credit Approval," *Jurnal Janitra Informatika dan Sistem Informasi*, vol. 1, no. 1, pp. 1–7, Apr. 2021, doi: 10.25008/janitra.v1i1.118.