

Penerapan *Explainable AI* dan Model *Stacking* Untuk Mengidentifikasi Faktor Risiko Stunting Balita

DOI: <http://dx.doi.org/10.35889/progresif.v22i1.3549>

Creative Commons License 4.0 (CC BY –NC)



Annisa Maulana Majid^{1*}, Ismasari Nawangsih²

Teknik Informatika, Universitas Pelita Bangsa, Bekasi, Indonesia

*e-mail *Corresponding Author* : annisa.maulanamajid@pelitabangsa.ac.id

Abstrack

Stunting remains a nutritional problem in children that is difficult to detect early due to its multifactorial causes and the limitations of conventional analytical approaches. This study aims to develop a stacking-based Machine Learning model for stunting prediction and to implement Explainable Artificial Intelligence (XAI) to interpret the risk factors influencing the prediction outcomes. The methods employed include Decision Tree, Random Forest, and Support Vector Machine (SVM) as base learners, and Logistic Regression as the meta-learner in an ensemble stacking framework, with performance evaluated using accuracy, precision, recall, and F1-score. The results indicate that the Decision Tree algorithm and the Stacking method achieved the best performance with 100% accuracy, while the XAI analysis identified current body weight, birth weight, and age as the primary factors influencing stunting prediction.

Keywords: *Stunting; Stacking; Machine Learning; Explainable Artificial Intelligence*

Abstrak

Stunting masih menjadi permasalahan gizi pada anak yang sulit dideteksi secara dini karena dipengaruhi oleh berbagai faktor dan keterbatasan analisis konvensional. Penelitian ini bertujuan mengembangkan model *stacking* berbasis *Machine Learning* untuk prediksi stunting serta menerapkan *Explainable Artificial Intelligence* (XAI) dalam menginterpretasikan faktor-faktor risiko yang berpengaruh terhadap hasil prediksi. Metode yang digunakan meliputi algoritma *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM) sebagai *base learner* dan *Logistic Regression* sebagai *meta learner* pada *ensemble stacking*, dengan evaluasi menggunakan *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa algoritma *Decision Tree* dan penerapan metode *Stacking* menghasilkan kinerja terbaik dengan tingkat akurasi 100%, sementara analisis XAI mengidentifikasi bahwa berat badan, BB lahir, dan usia merupakan faktor utama yang memengaruhi prediksi stunting.

Kata kunci: *Stunting; Stacking; Machine Learning; Explainable Artificial Intelligence*

1. Pendahuluan

Stunting masih menjadi permasalahan kesehatan di Indonesia, dengan prevalensi balita 21,5% pada 2023 yang menurun menjadi 19,8% pada 2024 [1]. Meskipun menurun, angka tersebut belum mencapai target nasional 14% sebagaimana ditetapkan dalam Rencana Pembangunan Jangka Menengah Nasional (RPJMN) [2]. Stunting memiliki dampak jangka panjang terhadap pertumbuhan fisik, perkembangan kognitif, kualitas pendidikan, serta produktivitas ekonomi di masa depan, sehingga pencegahan dan penanganannya menjadi isu yang sangat krusial.

Permasalahan stunting bersifat kompleks dan multidimensional. Faktor penyebabnya tidak hanya berkaitan dengan asupan gizi, tetapi juga kesehatan ibu, pola asuh, sanitasi lingkungan, serta kondisi sosial dan ekonomi keluarga. Kompleksitas interaksi antar faktor tersebut menyebabkan pendekatan analisis konvensional sering kali belum mampu menangkap pola hubungan yang tersembunyi secara optimal. Oleh karena itu, diperlukan pendekatan analitik berbasis teknologi yang mampu memprediksi risiko stunting secara akurat sekaligus

menghasilkan informasi yang mudah dipahami dan dapat ditindaklanjuti oleh tenaga kesehatan maupun pembuat kebijakan. Salah satu pendekatan yang banyak digunakan adalah *Machine Learning* karena kemampuannya dalam mengidentifikasi pola tersembunyi, memodelkan hubungan nonlinier, serta meningkatkan akurasi prediksi dibandingkan metode konvensional. Meskipun memiliki keunggulan tersebut, model *Machine Learning* umumnya bersifat *black box* sehingga proses pengambilan keputusannya sulit dipahami oleh pengguna non-teknis. Oleh karena itu, penelitian ini mengintegrasikan konsep *Explainable Artificial Intelligence* (XAI) guna meningkatkan transparansi hasil prediksi. *Explainable AI* (XAI) memiliki fungsi untuk meningkatkan transparansi dan interpretabilitas hasil prediksi, memberikan penjelasan hasil prediksi yang bersifat *black box* atau sulit dijelaskan [3]. Pendekatan XAI, seperti *Shapley Additive Explanations* (SHAP) dan *Local Interpretable Model-agnostic Explanations* (LIME), digunakan untuk menjelaskan kontribusi masing-masing variabel terhadap keputusan model. Dengan demikian, hasil prediksi tidak hanya bersifat akurat, tetapi juga memberikan pemahaman yang jelas mengenai faktor-faktor risiko dominan yang memengaruhi kejadian stunting.

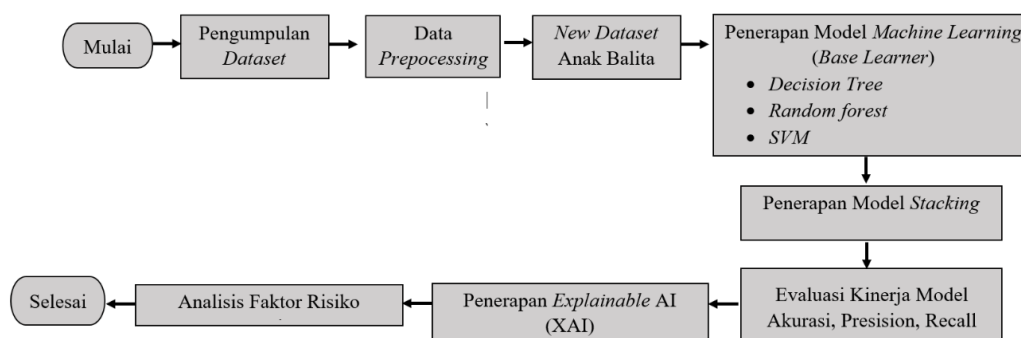
Penelitian-penelitian sebelumnya umumnya menggunakan algoritma tunggal, seperti *Naïve Bayes*, *Decision Tree*, *Support Vector Machine* (SVM), atau *Random Forest*, dengan tingkat akurasi yang bervariasi. Penelitian oleh La Denna Hasri Monsari, dkk tentang klasifikasi stunting dengan *dataset* berasal dari Kaggle.com berdasarkan usia, jenis kelamin, dan tinggi badan menggunakan algoritma *Decision Tree* menghasilkan akurasi sebesar 95,73% [4]. Penelitian oleh Nada Rizki Febriyanti, dkk tentang prediksi stunting balita dengan *dataset* berasal dari Dinas Kesehatan Kota Banjarmasin berdasarkan usia, berat badan, tinggi badan, dan z-score menghasilkan akurasi sebesar 92% pada algoritma SVM, akurasi sebesar 91% pada algoritma *Random Forest*, dan akurasi sebesar 92% pada algoritma *Logistic Regression* [5]. Penelitian oleh Rizqi Dwi Yuniarsyih, dkk tentang rasio kasus stunting dengan *dataset* berasal dari Kementerian Kesehatan RI tahun 2023 berdasarkan indeks ketahanan pangan, sanitasi, imunisasi dasar, rasio posyandu, pemberian ASI eksklusif, tempat pengelolaan makanan, dan lama sekolah menghasilkan akurasi sebesar 76% pada algoritma Logistik Biner dan 78% pada algoritma *Random Forest* [6]. Penelitian oleh M.U. Fido Millano, dkk dengan *dataset* berasal dari Kaggle.com berdasarkan usia, jenis kelamin, berat badan, dan tinggi badan menghasilkan akurasi sebesar 89,75% pada algoritma *Decision Tree* [7]. Penelitian oleh Mhd Adi Setiawan Aritonang, dkk dengan *dataset* berasal dari klasifikasi penyakit stunting berdasarkan karakteristik riwayat kesehatan, pola makan, status ekonomi keluarga, dan sanitasi rumah menghasilkan akurasi sebesar 74% pada algoritma *K-Nearest Neighbor* (K-NN) [8]. Meskipun beberapa penelitian menunjukkan hasil yang cukup tinggi, keterbatasan utama terletak pada rendahnya kemampuan interpretasi model terhadap faktor risiko stunting.

Penelitian ini berfokus pada upaya peningkatan akurasi prediksi risiko stunting pada balita dengan memanfaatkan pendekatan *Ensemble Machine Learning*. *Machine Learning* memungkinkan sistem melakukan pembelajaran dari data yang ada serta dapat mengidentifikasi pola-pola kompleks dan menghasilkan prediksi [9]. Algoritma *Machine Learning* (ML) menjadi salah satu solusi inovatif yang digunakan untuk mengatasi masalah kompleksitas klasifikasi [10]. Metode *ensemble* dapat digunakan untuk mengoptimasi algoritma klasifikasi dari *Machine Learning* untuk peningkatan performa, beberapa jenis *ensemble* antara lain: *bagging*, *boosting*, *voting*, *stacking*, dan *blending* [11]. *Ensemble Learning* dapat meningkatkan hasil prediksi dengan menggabungkan beberapa model yang unggul [12]. Data penelitian diperoleh dari catatan kesehatan balita, hasil pemantauan posyandu, serta survei pendukung yang relevan. Data tersebut melalui tahapan prapemrosesan sebelum digunakan dalam pelatihan model. Kinerja model *ensemble* kemudian dievaluasi menggunakan metrik akurasi, presisi, dan recall untuk menentukan metode terbaik dalam memprediksi stunting. Tujuan penelitian ini adalah menerapkan metode *Machine Learning* dan *Stacking* serta menggunakan *Explainable Artificial Intelligence* (XAI), guna meningkatkan akurasi prediksi sekaligus memberikan interpretasi yang transparan terhadap faktor-faktor risiko dominan stunting. Oleh karena itu, kontribusi utama penelitian ini adalah pengembangan model prediksi stunting yang mengombinasikan pendekatan *Ensemble Machine Learning* menggunakan *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM), *Logistic Regression*, serta *Stacking* dengan *Explainable Artificial Intelligence* (XAI). Hasil penelitian dapat mendukung pengambilan keputusan berbasis data yang lebih efektif oleh tenaga kesehatan dan pembuat kebijakan, serta dapat diimplementasikan secara langsung dalam upaya pencegahan stunting pada layanan kesehatan dasar seperti posyandu.

3. Metodologi

3.1 Tahapan Penelitian

Penelitian ini menerapkan metode *ensemble* yaitu *Stacking* menggunakan algoritma *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM) sebagai *base learner* atau model dasar, dan algoritma *Logistic Regression* sebagai *meta learner*. Proses penelitian ini dilakukan dengan menggunakan *python* pada platform *Google Collab*. Gambar 1 menunjukkan alur tahapan penelitian sebagai berikut:



Gambar 1. Tahapan Alur Penelitian

Tahap pertama dilakukan dengan pengumpulan *dataset* yang diambil dari posyandu tanjung XXIV Ciantra, kemudian dilakukan tahap data *preprocessing* untuk membersihkan data yang terdapat *noise* dan data yang tidak seimbang. Setelah dilakukan tahap *preprocessing* maka menghasilkan *dataset* baru yang siap dipakai untuk diimplementasikan menggunakan model tunggal algoritma *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM), lanjut pada tahap penerapan model *Stacking* menggunakan algoritma *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM) sebagai *base learner* atau model dasar, dan algoritma *Logistic Regression* sebagai *meta learner*. *Logistic Regression* dipilih karena mampu memproses hasil dari beberapa model dengan cepat dan efisien, struktur yang sederhana dan minim risiko *overfitting*, model ini cocok digunakan sebagai penggabung dalam sistem *stacking* tanpa menambah beban komputasi. Hasil pemodelan berupa *accuracy*, *presision*, *recall*, dan *F1-score*. Setelah tahap evaluasi model dilanjutkan pada penerapan *Explainable AI* (XAI) untuk menghasilkan analisa faktor risiko stunting pada anak balita.

3.2. Pengumpulan *Dataset*

Penelitian ini menggunakan *dataset* yang berasal dari data Posyandu Tanjung XXIV Ciantra yang berjumlah 130 data. *Dataset* balita tersebut terdiri dari 19 atribut yaitu NIK, Nama, Jenis Kelamin, Nama Orang Tua, Alamat, Tanggal Lahir, BB Lahir, TB Lahir, Usia Saat Ukur, Berat, Tinggi, Cara Ukur, BB/U, Zero Score BB/U, TB/U, Zero Score TB/U, BB/TB, Zero Score BB/TB, Naik Berat Badan dan 1 label yaitu Kategori Status Gizi. Label Kategori Status Gizi terdiri dari 6 kategori, sebagai berikut:

Tabel 1. Label Kategori Status Gizi

No.	Kategori Status Gizi	Jumlah Data
1	Underweight	1
2	Nomal	113
3	Stunting	12
4	Beresiko Gizi Lebih	2
5	Obesitas	1
6	Gizi Lebih	1
Total		130

3.3. Data Preprocessing

Tahapan *preprocessing* pada penelitian ini meliputi *data cleaning* (pembersihan data), reduksi data, transformasi data, pembentukan variabel target, serta pemisahan data menjadi variabel prediktor dan variabel target.

- 1) Pembersihan data (*data cleaning*) dilakukan dengan memeriksa adanya nilai kosong (*missing values*), kesalahan format data, dan inkonsistensi nilai pada setiap atribut. Hasil pemeriksaan menunjukkan bahwa *dataset* tidak mengandung nilai kosong, sehingga data dapat langsung digunakan pada tahap selanjutnya.
- 2) Reduksi data dengan menghapus kolom-kolom atribut yang tidak relevan terhadap tujuan penelitian. Tahap ini bertujuan untuk mengurangi kompleksitas data untuk meningkatkan hasil akurasi *dataset*. Kolom atribut yang dihapus meliputi NIK, Nama, Nama Orang Tua, Alamat, Tanggal Lahir, Cara Ukur, BB/U, Zero Score BB/U, TB/U, Zero Score TB/U, BB/TB, Zero Score BB/TB, dan Naik Berat Badan.
- 3) Transformasi dan pengkodean data untuk menyesuaikan format fitur, seperti jenis kelamin, dikonversi ke dalam bentuk numerik menggunakan teknik label encoding, dengan Laki-Laki = 0 dan Perempuan = 1.
- 4) Konversi satuan usia dari satuan bulan ke satuan tahun untuk menyederhanakan skala data dan memudahkan interpretasi hasil pemodelan. Sebagai contoh, seorang balita memiliki usia pengukuran sebesar 24 bulan, yang setelah dilakukan konversi menjadi 2 tahun dengan membagi nilai usia dalam bulan dengan angka 12.
- 5) Pembentukan variabel target dengan mengelompokkan status gizi menjadi kelas biner. Kategori status gizi yang menunjukkan kondisi berisiko, seperti stunting dan *underweight*, dikelompokkan sebagai kelas positif, sedangkan kategori gizi normal dan lebih dikelompokkan sebagai kelas negatif.

Tabel 2. Pembentukan Variabel Target

No.	Kategori Status Gizi	Jumlah Data	Kelas (y_binary)	Keterangan Label
1	<i>Underweight</i>	1	1	Positif
2	Nomal	113	0	Negatif
3	Stunting	12	1	Positif
4	Beresiko Gizi Lebih	2	0	Negatif
5	Obesitas	1	0	Negatif
6	Gizi Lebih	1	0	Negatif
Total		130		

- 6) Data dipisahkan menjadi variabel prediktor (X) dan variabel target (y) sebagai tahap akhir *preprocessing* untuk mempersiapkan data sebelum dilakukan pembagian data latih dan data uji serta proses pemodelan *Machine Learning*.

3.4. Teknik Pemodelan

Pemodelan dalam penelitian ini mengimplementasikan algoritma tunggal terlebih dahulu yaitu *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM), kemudian mengimplementasikan metode *Stacking* dengan *Decision Tree*, *Random Forest*, SVM sebagai *base learner* atau model dasar, dan algoritma *Logistic Regression* sebagai *meta learner*.

1) Decision Tree

Decision Tree merupakan metode *Machine Learning* untuk digunakan dengan aturan *if-then* dalam proses klasifikasi atau prediksi, yang digambarkan dalam bentuk struktur bercabang pohon yang menghasilkan keputusan akhir [13]. *Decision tree* adalah model yang direpresentasikan dalam bentuk struktur pohon, di mana setiap simpul internal digunakan untuk melakukan pengujian terhadap atribut, cabang merupakan hasil dari pengujian tersebut, dan simpul daun menunjukkan keputusan atau kelas akhir. Apabila variabel target bersifat diskret atau bernilai terbatas, maka model ini dikenal sebagai *classification tree*, dengan simpul daun berisi label kelas dan jalur cabang menggambarkan kombinasi atribut yang mengarah pada kelas tertentu [14]. *Decision Tree* merupakan pendekatan pemodelan yang menyusun data dalam struktur pohon bercabang sehingga mempermudah proses interpretasi dan implementasi,

meskipun model ini cenderung memiliki tingkat variasi yang relatif tinggi [15]. Rumus persamaan algoritma *Decision Tree* sebagai berikut:

Inisialisasi *dataset* untuk mempersiapkan atribut yang akan digunakan, kemudian menghitung *entropy* untuk menghitung ketidakpastian himpunan data [16]:

$$Entropy(S) = \sum_{i=0}^n -p_i * \log_2 p_i \dots \dots \dots (1)$$

Menghitung Information *Gain* untuk mengetahui besar pengurangan *entropy*

$$Gain(S, A) = Entropy(S) - \sum \frac{|S_v|}{|S|} * Entropy(S_v) \dots \dots \dots (2)$$

Menghitung *Gain Ratio* untuk mengetahui bias dalam atribut

$$GainRatio(A) = \frac{Gain(S, A)}{SplitInfo(A)} \dots \dots \dots (3)$$

$$SplitInfo(A) = - \sum \frac{|S_v|}{|S|} \log_2 \left(\frac{|S_v|}{|S|} \right) \dots \dots \dots (4)$$

Selanjutnya pemilihan atribut terbaik dengan *Gain Ratio* tertinggi sebagai simpul akar, pembuatan cabang daun berdasarkan nilai atribut, tahap akhir yaitu *pruning*/pemangkasan pohon untuk menghindari *overfitting*.

Keterangan:

- S : Himpunan Kasus
- N : Jumlah Pratisi S
- Pi : Proporsi dari Si ke S
- A : Atribut
- Sv : Jumlah kasus pada pratisi ke – v

2) Random Forest

Random Forest merupakan salah satu algoritma *Machine Learning* berbasis *Ensemble Learning* yang memanfaatkan kumpulan *Decision Tree* untuk menghasilkan prediksi yang lebih akurat dan stabil. Setiap pohon dikonstruksi menggunakan data dan subset fitur yang dipilih secara acak melalui teknik *bagging*, sehingga risiko *overfitting* dapat dikurangi dan kemampuan generalisasi model meningkat [17]. *Random Forest* mampu menghasilkan prediksi dengan tingkat akurasi yang tinggi disertai kebutuhan komputasi yang relatif efisien, serta menunjukkan ketahanan terhadap linearitas multivariant dan gangguan noise. Selain itu, algoritma ini efektif menangani data berdimensi tinggi tanpa memerlukan proses reduksi dimensi dan mampu memberikan informasi mengenai tingkat kepentingan setiap fitur dalam proses klasifikasi [18]. *Random Forest* merupakan algoritma *Ensemble Learning* untuk tugas klasifikasi dan regresi yang mengombinasikan prediksi dari banyak *Decision Tree*, di mana keputusan klasifikasi ditentukan berdasarkan kelas dengan jumlah prediksi terbanyak [19]. Rumus persamaan Random Forest sebagai berikut [20]:

$$f(x) = \sum_{i=1}^m w_i h_i(x) \dots \dots \dots (5)$$

$$F(x) = \frac{1}{M} \sum_{i=1}^M f_i(x) \dots \dots \dots (6)$$

Keterangan:

- f(x) : Output dari pohon keputusan
- m : Jumlah node dalam pohon keputusan
- w_i : Bobot dari setiap node dalam pohon keputusan
- h_i(x) : fungsi yang menghasilkan nilai 0 atau 1
- F(x) : Output dari *Random Forest*
- M : Jumlah pohon keputusan dalam Random Forest
- f_i(x) : Output dari pohon keputusan ke-1

3) Support Vector Machine (SVM)

SVM merupakan metode *Machine Learning* yang memiliki *hyperplane* optimal yang memisahkan dua kelas dengan margin maksimum, sehingga dapat memberikan generalisasi yang baik, terutama dataset yang berdimensi tinggi dan memiliki sampel terbatas [21]. Keunggulan SVM dapat mendeteksi *hyperplane* dengan optimal, tahan terhadap *noise*, menangani modulasi nonlinier, serta memiliki kemampuan generalisasi yang baik [22]. SVM mampu menangani permasalahan regresi karena memiliki kemampuan generalisasi baik. SVM dapat menentukan parameter serta fungsi kernel yang sesuai, sehingga dapat menangkap pola hubungan nonlinier meskipun jumlah data terbatas, dengan mencari solusi terbaik yang menyeimbangkan kesalahan prediksi dan kesederhanaan model melalui perhitungan matematika yang stabil [23]. Rumus persamaan SVM sebagai berikut [24]:

$$(w \cdot x_i) + b = 0 \dots\dots\dots (7)$$

$$(w \cdot x_i) + b \leq -1, y_1 = -1 \text{ dan } (w \cdot x_i) + b \geq 1, y_1 = 1 \dots\dots\dots (8)$$

Keterangan:

w : Vektor bobot yang menentukan arah *hyperplane* pemisah
 x_i : Vektor fitur ke- i inpur sample ke- i
 b : Bias yang menggeser posisi *hyperplane* terhadap titik pusat
 i : indeks data
 y_i : label kelas dari data ke- i

4) Logistic Regression

Logistic Regression salah satu metode statistik yang digunakan untuk memodelkan hubungan antara variabel dependen biner dengan satu atau lebih variabel independen [25]. *Logistic Regression* digunakan untuk memodelkan hubungan antara variabel bebas dan variabel terikat yang berbentuk kategori. *Logistic Regression* digunakan untuk mengestimasi probabilitas terjadinya suatu peristiwa dan dapat diterapkan baik pada kasus dengan dua kategori maupun lebih dari dua kategori [26]. *Logistic Regression* memodelkan peluang terjadinya suatu peristiwa dengan menghubungkan probabilitas kejadian tersebut dengan kombinasi variabel independen melalui fungsi logit yang mengubah probabilitas menjadi nilai numerik tanpa batas, sehingga mudah diolah oleh model statistic [27]. Rumus persamaan *Logistic Regression* sebagai berikut [28]:

$$\ln\left(\frac{p}{1-p}\right) = B_0 + B_1 X \dots\dots\dots (9)$$

$$p = \frac{e^{(B_0 + B_1 X)}}{(1 + e^{(B_0 + B_1 X)})} \dots\dots\dots (10)$$

Keterangan:

Ln : Logaritma natural
 p : Probabilitas kejadian
 B_0 : Konstanta intersep
 B_1 : Koefisien regresi variabel x
 X : Variabel independen
 e : Bilangan eksponensial

5) Stacking

Stacking merupakan strategi *ensemble* yang menggabungkan model dasar atau *base learner* menggunakan *meta learner*, sehingga mampu mengendalikan bias dan variansi secara bersamaan. *Stacking* dapat memberikan bobot berbeda pada tiap model dasar [29]. *Stacking ensemble* menggabungkan *base learner* dan *meta learner* untuk menghasilkan kinerja prediksi yang lebih baik. Pendekatan ini menggabungkan hasil keluaran dari setiap model dasar sehingga kelebihan masing-masing model dapat dimanfaatkan secara optimal [30]. *Stacking* adalah teknik *ensemble* yang menggabungkan prediksi beberapa model dasar melalui *meta-learner* untuk menghasilkan prediksi akhir dengan akurasi yang lebih baik [31]. Metode *Stacking* diterapkan dengan membangun beberapa model menggunakan algoritma yang beragam pada level dasar.

Prediksi dari model-model tersebut kemudian digunakan sebagai masukan bagi model pada level berikutnya untuk menghasilkan prediksi akhir. Dalam pendekatan ini, terdapat dua jenis model, yaitu model dasar (level-0) dan model meta (level-1), di mana model level-1 berperan mempelajari keluaran yang dihasilkan oleh model level-0 [32].

6) *Explainable Artificial Intelligence* (XAI)

Explainable Artificial Intelligence (XAI) merupakan pendekatan dalam kecerdasan buatan yang berfokus pada penyediaan penjelasan terhadap proses dan dasar pengambilan keputusan atau prediksi oleh model AI, sehingga hasilnya dapat dipahami dan dipercaya oleh manusia [33]. *Explainable Artificial Intelligence* (XAI) adalah pendekatan dalam kecerdasan buatan yang bertujuan membuat proses dan hasil keputusan model AI dapat dijelaskan dan dipahami oleh manusia, sehingga pengguna mengetahui mengapa dan bagaimana suatu prediksi dihasilkan serta dapat meningkatkan kepercayaan terhadap model tersebut [34]. Pengembangan *Explainable Artificial Intelligence* (XAI) dilakukan melalui tiga tahapan utama. Tahap awal menekankan penggunaan visualisasi dan simulasi sebagai sarana untuk membantu pengguna non-teknis memahami cara kerja sistem, sehingga kepercayaan dan interaksi dapat terbentuk. Tahap selanjutnya berfokus pada penyajian penjelasan yang lebih intuitif terhadap keputusan yang dihasilkan oleh AI, dengan memanfaatkan metode seperti SHAP dan LIME, yang berperan penting dalam sistem pembelajaran terawasi, misalnya pada deteksi penipuan keuangan. Pada tahap akhir, XAI diintegrasikan ke dalam alur proses bisnis secara menyeluruh, sehingga hasil keluaran AI dapat berinteraksi secara efektif dengan aturan bisnis dan mendukung penerapan lanjutan, seperti layanan pelanggan yang bersifat personal [35]. SHAP tidak hanya berkontribusi pada peningkatan akurasi model, tetapi juga menyediakan pemahaman yang dapat digunakan secara praktis melalui pengukuran pengaruh masing-masing fitur. Mengingat parameter model memiliki peran krusial dalam menentukan ketepatan prediksi [36].

3.5. Tahap Validasi dan Evaluasi

Pada tahap validasi, data dibagi menggunakan metode train–test split menjadi data latih (*training*) dan data uji (*testing*) untuk menguji kinerja model pada data yang tidak digunakan selama pelatihan. Pembagian data dengan proporsi 80% data latih dan 20% data uji. Data latih digunakan untuk membangun model klasifikasi, sedangkan data uji digunakan untuk memvalidasi kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya. Selanjutnya, evaluasi model dilakukan menggunakan *confusion matrix* untuk menggambarkan hasil klasifikasi, yang digunakan sebagai dasar perhitungan nilai *accuracy*, *precision*, *recall*, dan *F1-score*.

4. HASIL DAN PEMBAHASAN

4.1 Hasil Penelitian

Penelitian ini dilakukan dengan menerapkan model *Machine Learning*, model *Stacking*, dan XAI dengan menggunakan bahasa python pada google collab. Data dibagi menjadi 80% data latih sebesar 104 data dan 20% data uji sebesar 26 data dari total 130 data. Pembagian dilakukan secara acak menggunakan `random_state=42` dan teknik stratifikasi untuk menjaga keseimbangan distribusi kelas. Meskipun pembagian data dilakukan secara stratifikasi, jumlah data antar kelas tetap tidak seimbang. Oleh karena itu, pada tahap pemodelan digunakan parameter `class_weight='balanced'` untuk mengurangi bias terhadap kelas mayoritas. Pengujian pertama menerapkan algoritma *Decision Tree*. Berikut Gambar 2. menunjukkan sampel data hasil prediksi dari algoritma *Decision Tree*.

```
=== Sampel Hasil Prediksi Decision Tree (Balanced) ===
Aktual  Prediksi_DT
0      Normal      Normal
1      Normal      Normal
2      Normal      Normal
3      Normal      Normal
4      Normal      Normal
5      Normal      Normal
6      Stunting    Stunting
7      Stunting    Stunting
8      Normal      Normal
9      Normal      Normal
```

Gambar 2. Sampel Data Hasil Prediksi Algoritma *Decision Tree*

Berdasarkan sampel hasil prediksi, data berhasil diklasifikasikan dengan benar pada algoritma *Decision Tree*. Berikut Gambar 3. menunjukkan hasil evaluasi algoritma *Decision Tree*.

```

=== Decision Tree Balanced ===
1.0
[[23  0]
 [ 0  3]]

```

	precision	recall	f1-score	support
0	1.00	1.00	1.00	23
1	1.00	1.00	1.00	3
accuracy			1.00	26
macro avg	1.00	1.00	1.00	26
weighted avg	1.00	1.00	1.00	26

Gambar 3. Hasil Evaluasi Algoritma *Decision Tree*

Model kedua dengan menerapkan algoritma *Random Forest*. Berikut Gambar 4. menunjukkan sampel data hasil prediksi dari algoritma *Random Forest*.

```

=== Sampel Hasil Prediksi Random Forest (Balanced) ===
Aktual Prediksi_RF
0 Normal Normal
1 Normal Normal
2 Normal Normal
3 Normal Normal
4 Normal Normal
5 Normal Normal
6 Stunting Normal
7 Stunting Stunting
8 Normal Normal
9 Normal Normal

```

Gambar 4. Sampel Data Hasil Prediksi Algoritma *Random Forest*

Hasil prediksi *Random Forest* menunjukkan sebagian besar data terklasifikasi dengan benar, dengan satu kesalahan pada kelas *Stunting*. Secara umum, model mampu mendeteksi kelas *Stunting* dengan cukup baik. Berikut Gambar 5. menunjukkan hasil evaluasi algoritma *Random Forest*.

```

=== Random Forest Balanced ===
0.9615384615384616
[[23  0]
 [ 1  2]]

```

	precision	recall	f1-score	support
0	0.96	1.00	0.98	23
1	1.00	0.67	0.80	3
accuracy			0.96	26
macro avg	0.98	0.83	0.89	26
weighted avg	0.96	0.96	0.96	26

Gambar 5. Hasil Evaluasi Algoritma *Random Forest*

Model ketiga dengan menerapkan algoritma SVM. Berikut Gambar 6. menunjukkan sampel data hasil prediksi dari algoritma SVM.

```

=== Sampel Hasil Prediksi SVM (Balanced) ===
Aktual Prediksi_SVM
0 Normal Normal
1 Normal Stunting
2 Normal Normal
3 Normal Stunting
4 Normal Normal
5 Normal Normal
6 Stunting Normal
7 Stunting Normal
8 Normal Normal
9 Normal Stunting

```

Gambar 6. Sampel Data Hasil Prediksi Algoritma SVM

Berdasarkan sampel hasil prediksi, model SVM mampu mengklasifikasikan sebagian besar data kelas Normal dengan baik, namun masih mengalami kesalahan dalam mendeteksi kelas Stunting, yang menunjukkan bahwa sensitivitas model terhadap kasus Stunting masih perlu ditingkatkan. Berikut Gambar 7. menunjukkan hasil evaluasi algoritma SVM.

```

=== SVM Balanced ===
0.5769230769230769
[[14  9]
 [ 2  1]]

```

	precision	recall	f1-score	support
0	0.88	0.61	0.72	23
1	0.10	0.33	0.15	3
accuracy			0.58	26
macro avg	0.49	0.47	0.44	26
weighted avg	0.79	0.58	0.65	26

Gambar 7. Hasil Evaluasi Algoritma SVM

Model keempat dengan menggabungkan semua algoritma *Machine Learning* dengan menerapkan metode *Ensemble Stacking*. Algoritma *Decision Tree*, *Random Forest*, dan SVM digunakan sebagai *base learner* sedangkan *Logistic Regression* digunakan sebagai *meta learner*. Berikut Gambar 8. menunjukkan hasil dari penerapan metode *Stacking*.

```

=== STACKING ===

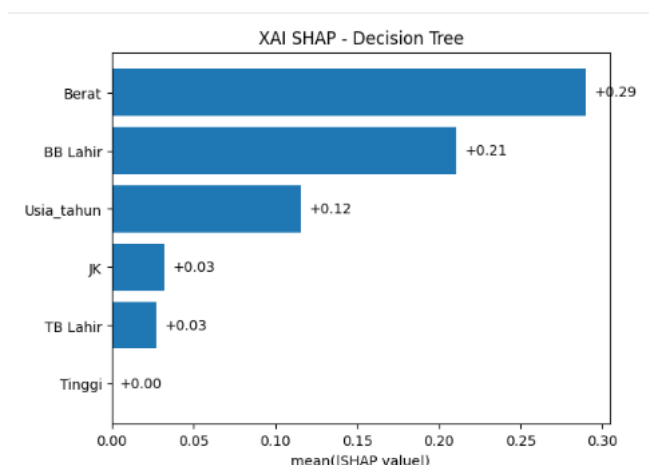
```

	precision	recall	f1-score	support
0	1.00	1.00	1.00	23
1	1.00	1.00	1.00	3
accuracy			1.00	26
macro avg	1.00	1.00	1.00	26
weighted avg	1.00	1.00	1.00	26

Gambar 8. Hasil Evaluasi Model *Stacking*

Model *Stacking* menghasilkan akurasi sebesar 81%. Berdasarkan *confusion matrix*, model mampu mendeteksi seluruh kasus stunting dengan hasil recall 100% tanpa adanya false negative. Hal ini menunjukkan model sangat baik dalam mendeteksi stunting. Namun demikian, precision pada kelas stunting masih rendah 38% karena terdapat beberapa data normal yang diklasifikasikan sebagai stunting. Dengan demikian, model cenderung sensitif terhadap kelas stunting dan meminimalkan risiko tidak terdeteksinya kasus stunting.

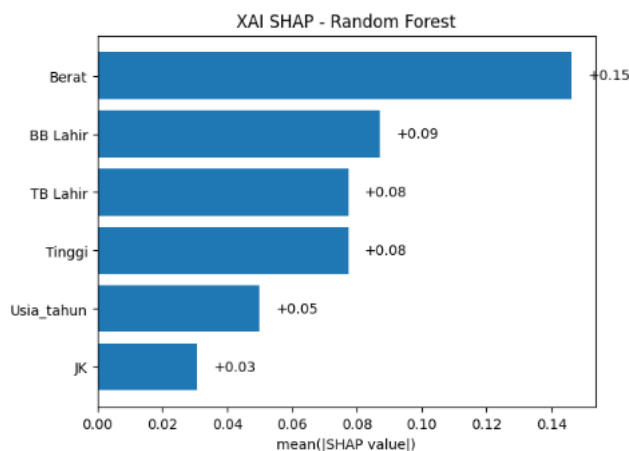
Selanjutnya penerapan model XAI pada setiap model. Model pertama menggunakan algoritma *Decision Tree*. Berikut Gambar 9 menunjukkan model XAI pada algoritma *Decision Tree*.



Gambar 9. Model XAI pada *Decision Tree*

Berdasarkan analisis SHAP pada model *Decision Tree*, variabel Berat memiliki kontribusi paling tinggi sebagai faktor penentu dalam proses prediksi stunting dengan nilai mean SHAP sebesar 0.29, diikuti oleh BB Lahir sebesar 0.21, Usia sebesar 0.12, Jenis Kelamin dan TB lahir masing-masing sebesar 0.03. Hal ini menunjukkan bahwa model *Decision Tree* sangat bergantung pada variabel berat badan sebagai indikator utama dalam pengambilan keputusan klasifikasi. Sementara itu, variabel Tinggi menunjukkan kontribusi yang sangat kecil terhadap prediksi model.

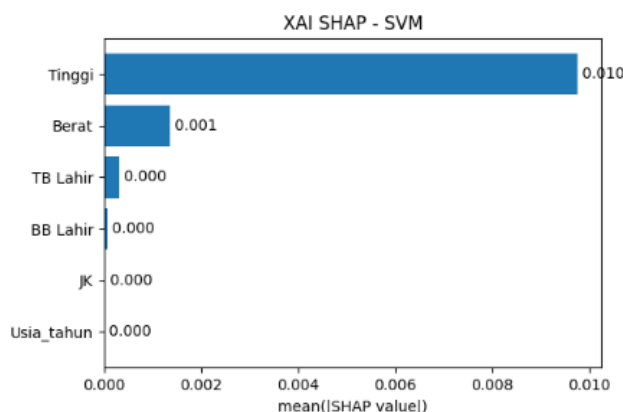
Selanjutnya model kedua penerapan XAI pada algoritma *Random Forest*. Berikut Gambar 10. menunjukkan hasil XAI pada algoritma *Random Forest*.



Gambar 10. Model XAI pada *Random Forest*

Berdasarkan analisis SHAP pada algoritma *Random Forest*, variabel Berat memiliki kontribusi paling tinggi sebagai faktor penentu dalam proses prediksi stunting dengan nilai mean SHAP sebesar 0,15. Hal ini menunjukkan bahwa variabel berat badan menjadi indikator utama yang paling memengaruhi keputusan prediksi model. Selanjutnya, variabel BB Lahir memiliki kontribusi sebesar 0,09, diikuti oleh TB Lahir dan Tinggi yang masing-masing memiliki nilai mean SHAP sebesar 0,08. Variabel Usia menunjukkan kontribusi sebesar 0,05, sedangkan Jenis Kelamin (JK) memiliki kontribusi paling kecil yaitu sebesar 0,03.

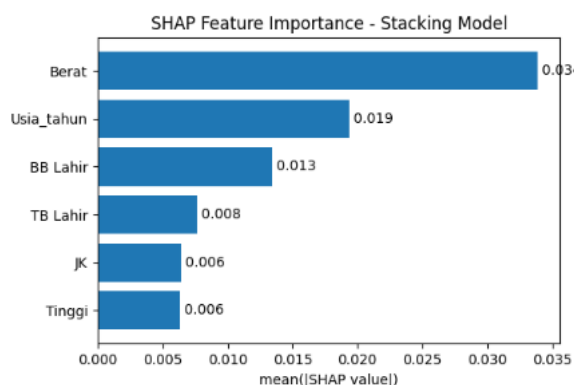
Model ketiga menerapkan model XAI pada algoritma SVM. Berikut Gambar 11. menunjukkan hasil XAI pada algoritma SVM.



Gambar 11. Model XAI pada SVM

Berdasarkan analisis SHAP pada algoritma SVM, variabel Tinggi memiliki kontribusi paling tinggi sebagai faktor penentu dalam prediksi stunting dengan nilai mean SHAP sebesar 0,010. Sementara itu, variabel Berat hanya memiliki kontribusi sebesar 0,001, dan variabel lain seperti TB Lahir, BB Lahir, JK, serta Usia menunjukkan nilai yang mendekati nol. Hal ini menunjukkan bahwa model SVM sangat bergantung pada variabel tinggi badan sebagai faktor utama dalam proses klasifikasi, sedangkan variabel lainnya memiliki pengaruh yang sangat kecil terhadap hasil prediksi.

Model keempat XAI pada algoritma *Stacking*. Berikut Gambar 12. menunjukkan hasil XAI pada algoritma *Stacking*.



Gambar 12. Model XAI pada Stacking

Berdasarkan analisis SHAP pada model Stacking, variabel Berat memiliki kontribusi paling tinggi sebagai faktor penentu dalam prediksi stunting dengan nilai mean SHAP sebesar 0,034. Selanjutnya, variabel Usia berkontribusi sebesar 0,019, diikuti oleh BB Lahir sebesar 0,013 dan TB Lahir sebesar 0,008. Sementara itu, variabel JK dan Tinggi memiliki kontribusi yang relatif kecil, masing-masing sebesar 0,006. Hal ini menunjukkan bahwa model Stacking tetap menjadikan berat badan sebagai faktor utama dalam klasifikasi, namun juga mempertimbangkan usia dan berat badan lahir sebagai variabel pendukung dalam proses prediksi.

4.2 Pembahasan

Penelitian menggunakan 3 model algoritma Machine Learning yaitu *Decision Tree*, *Random Forest*, dan SVM, kemudian menerapkan model *Ensemble Stacking* menggunakan algoritma *Decision Tree*, *Random Forest*, *Support Vector Machine* (SVM) sebagai *base learner* dan algoritma *Logistic Regression* sebagai *meta learner*. Model dievaluasi menghasilkan nilai *accuracy*, *precisiom*, *recall*, dan *F1-score*. Berikut tabel hasil evaluasi model:

Tabel 3. Hasil Evaluasi Model

Algoritma	Accuracy	Precision	Recall	F1-Score
Decision Tree	100%	100%	100%	100%
Random Forest	96%	98%	83%	89%
SVM	58%	49%	47%	44%
Stacking	100%	100%	100%	100%
DT+RF+RFM → LR)				

Berdasarkan hasil evaluasi performa model, algoritma *Decision Tree* dan *Stacking* (DT+RF+RFM → LR) menunjukkan kinerja terbaik dengan nilai *accuracy*, *precision*, *recall*, dan F1-score masing-masing sebesar 100%. Hal ini menunjukkan bahwa kedua model mampu mengklasifikasikan data dengan sangat baik pada dataset yang digunakan. Model *Random Forest* juga menunjukkan performa yang sangat baik dengan *accuracy* sebesar 96%, *precision* 98%, *recall* 83%, dan F1-score 89%. Meskipun tidak mencapai 100%, model ini tetap memiliki kemampuan klasifikasi yang tinggi dan relatif stabil. Sementara itu, SVM memiliki performa paling rendah dengan *accuracy* 58%, *precision* 49%, *recall* 47%, dan F1-score 44%, sehingga kurang optimal dalam memprediksi kasus stunting pada dataset penelitian ini. Secara keseluruhan, dapat disimpulkan bahwa model berbasis pohon (*Decision Tree* dan *Random Forest*), terutama yang dikombinasikan dalam metode *Stacking*, memberikan performa terbaik dalam klasifikasi stunting dibandingkan dengan model SVM.

Meskipun beberapa model menunjukkan akurasi tinggi, evaluasi performa belum menjelaskan proses pengambilan keputusan model. Oleh karena itu, dilakukan analisis XAI berbasis SHAP untuk mengidentifikasi kontribusi tiap variabel terhadap prediksi. Hasil analisis XAI menggunakan SHAP menunjukkan adanya perbedaan pola kontribusi fitur antar model. Model berbasis pohon (*Decision Tree* dan *Random Forest*) serta model ensemble (*Stacking*) secara konsisten menempatkan berat badan sebagai faktor paling dominan dalam prediksi stunting. Selain itu, berat badan lahir dan usia juga menunjukkan kontribusi yang cukup signifikan, terutama pada model ensemble yang memperlihatkan distribusi pengaruh fitur yang lebih seimbang. Sebaliknya, model SVM menunjukkan ketergantungan yang tinggi pada satu variabel utama, yaitu tinggi badan, dengan kontribusi fitur lain yang sangat minimal. Hal ini mengindikasikan bahwa SVM kurang mampu menangkap hubungan kompleks antar variabel dibandingkan model berbasis pohon. Secara keseluruhan, interpretasi SHAP memperlihatkan bahwa indikator antropometri, khususnya berat badan saat ini dan riwayat berat badan lahir, merupakan faktor yang paling konsisten berperan dalam pembentukan keputusan model terhadap risiko stunting.

Jika dibandingkan dengan penelitian terdahulu, temuan ini memperkuat hasil yang diperoleh oleh La Denna Hasri Monsari, dkk [4] dan M.U. Fido Millano, dkk [7] yang menunjukkan bahwa algoritma *Decision Tree* memiliki performa tinggi dalam klasifikasi stunting. Selain itu, hasil *Random Forest* yang mencapai akurasi 96% juga sejalan dengan temuan Nada Rizki Febriyanti, dkk [5] serta Rizqi Dwi Yuniarsyih, dkk [6] yang menyatakan bahwa model berbasis pohon cenderung memberikan hasil yang stabil dibandingkan algoritma lainnya. Namun, penelitian ini memberikan kontribusi tambahan melalui penerapan metode *Ensemble Stacking* yang mampu mempertahankan performa optimal serta integrasi analisis XAI berbasis SHAP. Dengan demikian, penelitian ini tidak hanya menguatkan temuan sebelumnya terkait efektivitas model berbasis pohon, tetapi juga memperluasnya melalui pendekatan interpretabilitas model sehingga mendukung pengembangan sistem prediksi stunting yang lebih akurat dan transparan dalam bidang kesehatan masyarakat.

5. Simpulan

Penelitian ini menunjukkan bahwa model berbasis pohon, khususnya *Decision Tree* dan *Stacking* (DT+RF+RFM → LR), memberikan performa terbaik dengan nilai *accuracy*, *precision*, *recall*, dan F1-score masing-masing sebesar 100%, diikuti oleh *Random Forest* dengan *accuracy* 96%, *precision* 98%, *recall* 83%, dan F1-score 89%, sementara SVM menunjukkan performa terendah dengan *accuracy* 58%, *precision* 49%, *recall* 47%, dan F1-score 44%. Analisis XAI berbasis SHAP mengungkapkan bahwa berat badan merupakan faktor paling dominan dalam prediksi stunting (mean SHAP hingga 0,29 pada *Decision Tree* dan 0,15 pada *Random Forest*), diikuti oleh berat badan lahir (hingga 0,21) dan usia (hingga 0,12). Temuan ini menegaskan bahwa indikator antropometri, khususnya berat badan saat ini dan riwayat berat badan lahir,

menjadi penentu utama risiko stunting, serta menunjukkan bahwa integrasi machine learning dan pendekatan interpretabel mampu memberikan prediksi yang akurat sekaligus transparan dalam mendukung pengambilan keputusan di bidang kesehatan masyarakat.

Penelitian selanjutnya disarankan untuk menggunakan *dataset* dengan jumlah sampel yang lebih besar dan distribusi kelas yang lebih seimbang guna menguji stabilitas model, terutama karena beberapa algoritma menunjukkan akurasi 100% yang berpotensi mengindikasikan *overfitting*. Selain itu, penambahan variabel lain seperti faktor sosial ekonomi, pola asuh, riwayat penyakit infeksi, sanitasi, dan asupan gizi harian perlu dipertimbangkan untuk memperkaya analisis faktor risiko stunting. Pengembangan metode juga dapat dilakukan dengan mengeksplorasi algoritma lain atau teknik validasi yang lebih ketat serta memperluas penerapan *Explainable AI* (XAI) agar interpretasi model semakin komprehensif dan dapat mendukung pengambilan kebijakan kesehatan yang lebih tepat sasaran.

Daftar Referensi

- [1] Kementerian Kesehatan Republik Indonesia., *Survei Status Gizi Indonesia (SSGI) 2024 dalam angka*. Jakarta: Badan Kebijakan Pembangunan Kesehatan, 2025.
- [2] P. A. K. dan Kinerja, *Policy Brief Policy Brief*, vol. 6, Oktober. Jakarta: Kementerian PPN/Bappenas, 2023.
- [3] S. Muliani, B. Sukma Negara, M. Irsyad, I. Iskandar, and J. Teknik Informatika, "Application of Shapley Additive Explanations (SHAP) in Deep Learning for Lung Disease Detection Using X-ray Images," *J. Artif. Intell. Softw. Eng.*, vol. 5, no. 2, pp. 709–719, 2025, doi: 10.30811/jaise.v5i2.7044.
- [4] L. Denna, H. Monasari, and N. Pratiwi, "Klasifikasi Stunting Pada Balita Menggunakan Algoritma Decision Tree Pada Machine Learning," vol. 5, no. 2, pp. 416–428, 2025.
- [5] N. R. Febriyanti, K. Kusriani, and A. D. Hartanto, "Analisis Perbandingan Algoritma SVM, Random Forest dan Logistic Regression untuk Prediksi Stunting Balita," *Edumatic J. Pendidik. Inform.*, vol. 9, no. 1, pp. 149–158, 2025, doi: 10.29408/edumatic.v9i1.29407.
- [6] R. D. Yuniarsih, "Analisis Regresi Logistik Biner dan Random Forest untuk Prediksi Faktor-Faktor Stunting di Pulau Jawa," *Euler J. Ilm. Mat. Sains Dan Teknol. ?*, vol. 13, no. 2, pp. 147–156, 2025, [Online]. Available: <https://ejurnal.ung.ac.id/index.php/Euler/article/view/31680/10964>
- [7] M. U. F. M. P, S. Kurniawan, and R. Gustriansyah, "KOMPUTA : Jurnal Ilmiah Komputer dan Informatika Penerapan Metode Decision Tree Dalam Klasifikasi Status Gizi Balita Application of the Decision Tree Method for Classifying the Nutritional Status of Toddlers," vol. 14, no. 2, pp. 2–10, 2025, doi: 10.34010/komputa.v14i2.
- [8] M. Adi, S. Aritonang, M. A. Siregar, S. A. Arnomo, and F. Farasalsabila, "Penerapan Algoritma Metaheuristik untuk Optimasi Nilai K pada KNN Pengelompokan Faktor Stunting," vol. 5, no. 3, pp. 611–618, 2025.
- [9] K. Sakthivel and S. A. Alexander, "An extensive critique on machine learning techniques for fault tolerance and power quality improvement in multilevel inverters," *Energy Reports*, vol. 12, August, pp. 5814–5833, 2024, doi: 10.1016/j.egyr.2024.11.016.
- [10] A. Wibowo and A. R. Isnain, "Implementasi Algoritma Machine Learning untuk Klasifikasi Suara Lingkungan," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 5, no. 2, pp. 616–625, 2025, doi: 10.57152/malcom.v5i2.1712.
- [11] N. K. Majid, C. Supriyanto, and A. Marjuni, "Peningkatan Keberagaman Data untuk Klasifikasi Penyakit Diabetes Berbasis Stacking Ensemble Learning," *J. Inform. J. Pengemb. IT*, vol. 10, no. 1, pp. 1–10, 2025, doi: 10.30591/jpit.v10i1.7375.
- [12] G. A. Cynthia Nur Azzahra, Yulison Herry Chrisnanto, "Klasifikasi Cuaca Jawa Barat menggunakan Ensemble Learning pada Data Meteorologi Weather Classification in West Java using Ensemble Learning on," vol. 14, no. 5, pp. 2028–2044, 2025.
- [13] D. M. Rodríguez, M. P. Cuéllar, and D. P. Morales, "On the fusion of soft-decision-trees and concept-based models," *Appl. Soft Comput.*, vol. 160, February, p. 111632, 2024, doi: 10.1016/j.asoc.2024.111632.
- [14] M. D. Lagzi, S. M. Sajadi, and M. Taghizadeh-Yazdi, "A hybrid stochastic data envelopment analysis and decision tree for performance prediction in retail industry," *J. Retail. Consum. Serv.*, vol. 80, February, p. 103908, 2024, doi: 10.1016/j.jretconser.2024.103908.
- [15] A. F. Azmi and A. Voutama, "Prediksi Churn Nasabah Bank Menggunakan Klasifikasi Random Forest Dan Decision Tree Dengan Evaluasi Confusion Matrix," *Komputa J. Ilm.*

- Komput. dan Inform.*, vol. 13, no. 1, pp. 111–119, 2024, doi: 10.34010/komputa.v13i1.12639.
- [16] S. N. Sari, P. Annisa, A. Nisa, D. Rahma, R. Prameswara Ritonga, and D. P. Utomo, "Teknik Data Mining Dengan Menggunakan Algoritma Decision Tree Untuk Mengetahui Pola Pemahaman Mahasiswa Pada Matakuliah Pemrograman," *Bull. Inf. Technol.*, vol. 6, no. 4, pp. 417–431, 2025, doi: 10.47065/bit.v5i2.2339.
- [17] G. Grimmer, R. Wenger, G. Forestier, and V. Chardon, "Automatic detection of in-stream river wood from random forest machine learning and exogenous indices using very high-resolution aerial imagery," *Environ. Model. Softw.*, vol. 190, February, p. 106460, 2025, doi: 10.1016/j.envsoft.2025.106460.
- [18] J. Zhang, "Impact of an improved random forest-based financial management model on the effectiveness of corporate sustainability decisions," *Syst. Soft Comput.*, vol. 6, February, p. 200102, 2024, doi: 10.1016/j.sasc.2024.200102.
- [19] S. T. Hamidou and A. Mehdi, "Enhancing IDS performance through a comparative analysis of Random Forest, XGBoost, and Deep Neural Networks," *Mach. Learn. with Appl.*, vol. 22, May, p. 100738, 2025, doi: 10.1016/j.mlwa.2025.100738.
- [20] J. P. Anggraini, Chaya Gladys Zhafirah A, and A. Desiani, "Perbandingan Algoritma Random Forest dan Extreme Gradient Boosting (XGBoost) dalam Klasifikasi Penyakit Gagal Jantung," *Komputika J. Sist. Komput.*, vol. 14, no. 2, pp. 149–157, 2025, doi: 10.34010/komputika.v14i2.16618.
- [21] S. Zhao, X. Liang, L. Wang, H. Zhang, G. Li, and J. Chen, "A fault diagnosis method for analog circuits based on EEMD-PSO-SVM," *Heliyon*, vol. 10, no. 18, p. e38064, 2024, doi: 10.1016/j.heliyon.2024.e38064.
- [22] A. Kumar, N. Gaur, and A. Nanthamornphong, "Machine learning RNNs, SVM and NN Algorithm for Massive-MIMO-OTFS 6G Waveform with Rician and Rayleigh channel," *Egypt. Informatics J.*, vol. 27, July, p. 100531, 2024, doi: 10.1016/j.eij.2024.100531.
- [23] W. An, B. Gao, J. Liu, J. Ni, and J. Liu, "Predicting hourly heating load in residential buildings using a hybrid SSA–CNN–SVM approach," *Case Stud. Therm. Eng.*, vol. 59, May, 2024, doi: 10.1016/j.csite.2024.104516.
- [24] D. N. Ayu Ofta Sari, Muhammad Iqbal, "Analisis Faktor Demografi Dan Sosial Ekonomi Untuk Mendeteksi Dini Risiko Putus Kuliah Menggunakan Metode Support Vector Machine (Svm) Dan Decision Tree (Studi Kasus : STMIK Triguna Dharma)," *J. Multimed*, vol. 07, no. 02, p. 17, 2025, [Online]. Available: <http://www.journal.cattleyadf.org/index.php/jatilima/article/view/1439%0Ahttp://www.journal.cattleyadf.org/index.php/jatilima/article/download/1439/797>
- [25] G. Zeng, "A comprehensive study of coefficient signs in weighted logistic regression," *Heliyon*, vol. 10, no. 15, p. e35040, 2024, doi: 10.1016/j.heliyon.2024.e35040.
- [26] T. M. Saunders, N. Cohn, and T. Hesser, "Insights into nearshore sandbar dynamics through process-based numerical and logistic regression modeling," *Coast. Eng.*, vol. 192, June, p. 104558, 2024, doi: 10.1016/j.coastaleng.2024.104558.
- [27] A. Balboa, A. Cuesta, J. González-Villa, G. Ortiz, and D. Alvear, "Logistic regression vs machine learning to predict evacuation decisions in fire alarm situations," *Saf. Sci.*, vol. 174, February, 2024, doi: 10.1016/j.ssci.2024.106485.
- [28] F. D. Pramakrisna, F. D. Adhinata, and N. A. F. Tanjung, "Aplikasi Klasifikasi SMS Berbasis Web Menggunakan Algoritma Logistic Regression," *Teknika*, vol. 11, no. 2, pp. 90–97, 2022, doi: 10.34148/teknika.v11i2.466.
- [29] Y. Zhang, C. Chen, F. Guo, and H. Wang, "Geographical origin identification of dendrobium officinale based on NNRW-stacking ensembles," *Mach. Learn. with Appl.*, vol. 18, September, p. 100594, 2024, doi: 10.1016/j.mlwa.2024.100594.
- [30] M. Rahman, N. Kamal, and N. F. Abdullah, "EDT-STACK: A stacking ensemble-based decision trees algorithm for tire tread depth condition classification," *Results Eng.*, vol. 22, May, p. 102218, 2024, doi: 10.1016/j.rineng.2024.102218.
- [31] M. Dostmohammadi, M. Z. Pedram, S. Hoseinzadeh, and D. A. Garcia, "A GA-stacking ensemble approach for forecasting energy consumption in a smart household: A comparative study of ensemble methods," *J. Environ. Manage.*, vol. 364, June, p. 121264, 2024, doi: 10.1016/j.jenvman.2024.121264.
- [32] N. O. Syamsiah and I. Purwandani, "Penerapan Metode Stacking Untuk Meningkatkan Akurasi Hasil Peramalan Konsumsi Listrik," *J. Syst. Comput. Eng.*, vol. 4, no. 1, pp. 15–25, 2023, doi: 10.47650/jsce.v4i1.665.

- [33] L. P. Di Bonito, L. Campanile, M. Iacono, and F. Di Natale, "An eXplainable Artificial Intelligence framework to predict marine scrubbers performances," *Eng. Appl. Artif. Intell.*, vol. 160, December, p. 111860, 2025, doi: 10.1016/j.engappai.2025.111860.
- [34] I. Vega-Rebolledo, A. J. Sánchez-García, J. J. Muñoz León, J. O. Ocharán-Hernández, and K. Cortés-Verdín, "Applying survival analysis and Explainable Artificial Intelligence to understand academic success in software engineering students," *Array*, vol. 28, July, p. 100540, 2025, doi: 10.1016/j.array.2025.100540.
- [35] A. Mohamed, K. Abdelqader, and K. Shaalan, "Explainable Artificial Intelligence: A systematic Review of Progress and Challenges," *Intell. Syst. with Appl.*, vol. 28, August, p. 200595, 2025, doi: 10.1016/j.iswa.2025.200595.
- [36] M. Chen, E. Zhao, L. Zhang, S. Zhang, and Y. Li, "Interpretable stacking-SHAP prediction model for dam deformation during long term operation," *Mech. Syst. Signal Process.*, vol. 242, p. 236, September 2025, 2026, doi: 10.1016/j.ymssp.2025.113646.