

KLASIFIKASI TUMBUH KEMBANG BALITA BERDASARKAN KPSP MENGGUNAKAN RANDOM FOREST PADA SISTEM BERBASIS WEB

Muhammad Daffa Ariq Fadilah^{1*}, Rini Mayasari², Carudin³, Purwantoro⁴, Iqbal Maulana⁵

^{1,2,3,4,5}Program Studi Informatika, Fakultas Ilmu Komputer, Universitas Singaperbangsa Karawang; Jl. HS. Ronggo Waluyo, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat 41361

Keywords: KPSP; Random Forest; SHAP; Explainable AI; Tumbuh Kembang Balita

Correspondent Email:

2210631170083@student.unsika.ac.id

Abstrak. Penyimpangan tumbuh kembang balita yang tidak terdeteksi dini berdampak permanen terhadap kualitas hidup anak. Kuesioner Pra Skrining Perkembangan (KPSP) merupakan instrumen standar deteksi dini, namun pelaksanaannya di lapangan masih manual sehingga data skrining tidak dimanfaatkan secara longitudinal, tidak teragregasi antar fasilitas, dan interpretasinya terbatas pada ambang skor tanpa penjelasan. Penelitian ini mengembangkan sistem berbasis web yang mengaugmentasi interpretasi KPSP menggunakan algoritma Random Forest yang dilengkapi penjelasan keputusan melalui metode Shapley Additive exPlanations (SHAP). Model dilatih menggunakan data sintesis berbasis aturan Permenkes 66/2014 sebanyak 4.800 baris dan divalidasi pada 105 data primer skrining kelompok usia enam bulan dari Puskesmas Ciampel. Konfigurasi hyperparameter ditetapkan berdasarkan praktik yang mapan dan diverifikasi melalui grid search dengan validasi silang lima lipatan. Model mencapai akurasi 0,8827 dan F1-macro 0,8580 pada data uji, dengan seluruh kesalahan klasifikasi hanya bergeser satu tingkat ke kelas yang bersebelahan; pada data primer, prediksi model konsisten dengan penilaian tenaga kesehatan. Sistem menyajikan hasil dalam tiga lapisan keluaran, yaitu hasil per-skrining beserta penjelasan SHAP, trajektori longitudinal per anak, dan rekap agregat tingkat fasilitas. Hasil menunjukkan bahwa augmentasi machine learning dapat memperkaya interpretasi KPSP tanpa menggantikan peran instrumen maupun tenaga kesehatan.



Copyright © 2026 The Author(s), Published by Jurnal Informatika dan Teknik Elektro Terapan. This is an open-access article distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (CC BY-NC 4.0).

Abstract. Undetected developmental disorders in toddlers have a permanent impact on a child's quality of life. The Pre-Screening Developmental Questionnaire (KPSP) is the standard early detection instrument, but its field implementation is still manual, so screening data is not used longitudinally, is not aggregated across facilities, and its interpretation is limited to score thresholds without explanation. This study develops a web-based system that augments KPSP interpretation using the Random Forest algorithm, equipped with decision explanations via the Shapley Additive exPlanations (SHAP) method. The model was trained on 4,800 rows of synthetic data derived from the Permenkes 66/2014 rules and validated on 105 primary screening records of the six-month age group from Puskesmas Ciampel. The hyperparameter configuration was set based on established practice and verified through grid search with five-fold cross-validation. The model achieved 0.8827 accuracy and 0.8580 macro-F1 on the test set, with all misclassifications shifting only one level to an adjacent class; on the primary data, model predictions were consistent with health workers' assessments. The system presents results in three output layers: per-screening results with SHAP explanations, longitudinal per-child trajectories, and facility-level aggregates. The results

show that machine learning augmentation can enrich KPSP interpretation without replacing the instrument or the role of health workers.

1. PENDAHULUAN

Masa balita merupakan periode emas tumbuh kembang manusia sehingga penyimpangan perkembangan yang tidak terdeteksi sejak dini berpotensi berdampak permanen terhadap kualitas hidup anak. World Health Organization mencatat prevalensi penyimpangan perkembangan pada balita Indonesia sebesar 7,51% [1], Kementerian Kesehatan melaporkan 5.530 kasus gangguan perkembangan anak pada periode 2020–2021 [2], dan Riset Kesehatan Dasar mencatat sekitar 16% balita mengalami gangguan perkembangan motorik [3]. Laporan global terbaru menyoroti kesenjangan deteksi dini dan akses intervensi bagi anak dengan keterlambatan perkembangan [4]. Sebagai instrumen standar deteksi dini untuk anak usia 3–72 bulan, pemerintah menetapkan Kuesioner Pra Skrining Perkembangan (KPSP) melalui Peraturan Menteri Kesehatan Nomor 66 Tahun 2014, yang menginterpretasikan jumlah jawaban “Ya” atas sepuluh pertanyaan menjadi tiga kategori: Sesuai, Meragukan, dan Penyimpangan [5].

Pada praktiknya, pelaksanaan skrining KPSP di lapangan masih dilakukan secara manual menggunakan formulir kertas, sehingga data hasil skrining praktis berhenti dimanfaatkan setelah formulir selesai diisi. Tidak terdapat pencatatan elektronik longitudinal untuk memantau trajektori perkembangan tiap anak, tidak terdapat agregasi otomatis antar-Posyandu sebagai bahan pengambilan keputusan tingkat fasilitas, dan orang tua tidak memiliki akses terhadap riwayat skrining anaknya. Selain itu, interpretasi berbasis ambang skor tidak menangkap interaksi antar-jawaban maupun informasi kontekstual pasien seperti riwayat kelahiran prematur, dan tidak menyediakan penjelasan atas hasil yang diberikan.

Penerapan machine learning pada berbagai tugas klasifikasi telah banyak dilaporkan, termasuk pada data berbahasa Indonesia [22]; pada domain tumbuh kembang balita, penerapannya umumnya dilakukan di luar instrumen KPSP dan tanpa integrasi ke dalam

sistem yang digunakan tenaga kesehatan. Bimawan dkk. membandingkan empat algoritma dan menempatkan Random Forest sebagai algoritma terbaik untuk deteksi stunting [7], sedangkan Majid dan Nawangsih menerapkan explainable AI pada faktor risiko stunting tanpa integrasi sistem [8]. Di sisi lain, digitalisasi KPSP oleh Apriningrum dkk. berfokus pada penghitungan skor berbasis aturan tanpa kemampuan klasifikasi maupun penjelasan [6]. Dengan demikian, belum terdapat penelitian yang menerapkan Random Forest dengan dukungan explainability pada instrumen KPSP dan mengintegrasikannya ke dalam sistem berbasis web.

Penelitian ini mengembangkan sistem berbasis web yang mengaugmentasi interpretasi hasil skrining KPSP menggunakan algoritma Random Forest [9] yang dilengkapi penjelasan keputusan melalui metode Shapley Additive exPlanations (SHAP) [10]. Sistem diposisikan sebagai alat bantu interpretasi—bukan pengganti instrumen KPSP maupun peran tenaga kesehatan—dan menyajikan hasil dalam tiga lapisan keluaran, yaitu hasil per-skrining individu beserta penjelasan SHAP, trajektori longitudinal per anak, dan rekap agregat tingkat fasilitas.

2. TINJAUAN PUSTAKA

KPSP merupakan instrumen skrining perkembangan yang terdiri atas sepuluh pertanyaan per kelompok umur, mencakup empat aspek perkembangan: gerak kasar, gerak halus, bicara-bahasa, dan sosial-kemandirian. Interpretasi hasil ditentukan dari jumlah jawaban “Ya”, yaitu Sesuai (9–10), Meragukan (7–8), dan Penyimpangan (≤ 6) [5].

Random Forest adalah algoritma ensemble yang membangun banyak pohon keputusan pada subset data dan fitur yang diacak, kemudian menggabungkan prediksinya melalui pemungutan suara mayoritas [9]. Algoritma ini unggul pada data tabular berdimensi sedang, tahan terhadap overfitting, dan mampu menangkap interaksi non-linear antar-fitur; perilaku serta strategi penyetelan hyperparameter-nya telah dikaji secara empiris [16].

SHAP (Shapley Additive exPlanations) adalah metode explainable AI yang menguraikan prediksi model menjadi kontribusi aditif tiap fitur berdasarkan konsep nilai Shapley dari teori permainan [10]. Sifat aditivitasnya menjamin bahwa penjumlahan nilai dasar dan seluruh kontribusi fitur sama dengan keluaran model, sehingga setiap prediksi dapat ditelusuri secara individual.

Pengembangan model mengikuti kerangka CRISP-DM yang mencakup fase Business Understanding, Data Understanding, Data Preparation, Modeling, Evaluation, dan Deployment [11]. Penelitian terdahulu yang relevan antara lain penerapan decision tree untuk monitoring tumbuh kembang di PAUD [19] dan prediksi kesiapan sekolah berbasis machine learning [20], yang menegaskan potensi klasifikasi pada domain perkembangan anak namun belum menyentuh integrasi explainability ke dalam instrumen KPSP.

3. METODE PENELITIAN

Penelitian mengikuti kerangka CRISP-DM [11] dengan enam fase yang diuraikan sebagai berikut. Seluruh alur bersifat reproducible dan diimplementasikan menggunakan pustaka scikit-learn.

1. Data dan Sumber Data

Penelitian menggunakan dua sumber data. Sumber pertama adalah data sintetis yang dibangun berdasarkan struktur instrumen KPSP dan aturan penilaian Permenkes 66/2014; penggunaan data sintetis merupakan pendekatan yang lazim pada domain kesehatan ketika data nyata terbatas [12]. Sumber kedua adalah 105 data primer hasil skrining KPSP kelompok usia enam bulan yang dikumpulkan melalui pelaksanaan Pemeriksaan Kesehatan Anak Terintegrasi (PKAT) di Puskesmas Ciampel, yang berperan sebagai data uji eksternal. Data sintetis dibangun berdasarkan sel, yaitu 100 sampel untuk setiap kombinasi 16 kelompok umur KPSP dan 3 kelas, menghasilkan 4.800 baris data latih (1.600 per kelas) dan 1.500 baris data uji internal; ukuran ini berada di atas ambang kecukupan data untuk teknik non-linear [13].

2. Persiapan Data

Tahap persiapan data mencakup pembersihan, transformasi fitur, dan

penyusunan struktur masukan. Fitur kategorikal dikodekan menjadi numerik dan sepuluh jawaban KPSP dipertahankan dalam bentuk biner. Label akhir ditetapkan melalui mekanisme penyesuaian (label drift) berbasis sejumlah kondisi klinis-kontekstual—antara lain kelemahan pada aspek bicara-bahasa pada usia tertentu, kelemahan pada lebih dari satu aspek perkembangan, faktor prematuritas, dan variasi penilaian klinis—agar label tidak semata mengikuti ambang skor, dengan proporsi perubahan terukur 14,33% pada data latih dan 15,13% pada data uji. Seluruh fitur disusun mengikuti feature contract berisi lima belas fitur berurutan tetap, yaitu sepuluh jawaban KPSP ditambah usia (bulan), kelompok umur, jenis kelamin, status prematur, dan usia gestasi, yang digunakan identik pada pelatihan maupun inferensi. Normalisasi tidak diterapkan karena Random Forest tidak sensitif terhadap skala fitur. Data latih diseimbangkan antar-kelas [14], sedangkan proporsi latih terhadap uji (76:24) mengikuti konvensi yang umum digunakan [15].

3. Pemodelan dan Evaluasi

Model Random Forest dilatih menggunakan label akhir hasil label drift. Konfigurasi hyperparameter ditetapkan melalui dua tahap: penetapan awal berdasarkan praktik yang mapan pada pemodelan data tabular dan penalaran regularisasi, kemudian verifikasi menyeluruh menggunakan grid search terhadap 288 kombinasi dengan validasi silang lima lipatan dan F1-macro sebagai metrik; pemilihan lima lipatan menyeimbangkan bias dan varians estimasi [15]. Grid search digunakan untuk menguji apakah konfigurasi awal telah mendekati konfigurasi optimal, bukan sebagai mekanisme pemilihan model. Konfigurasi awal dipertahankan berdasarkan kriteria one-standard-error, yaitu mempertahankan konfigurasi yang lebih sederhana apabila skornya masih berada dalam rentang satu simpangan baku dari skor terbaik [18]. Konfigurasi final menggunakan 200 pohon, kedalaman maksimum 12, min_samples_split 5, min_samples_leaf 2, dan class_weight seimbang; jumlah pohon yang memadai konsisten dengan temuan empiris bahwa performa Random Forest stabil setelah kisaran 64–128 pohon [17]. Evaluasi dilakukan dengan menghitung confusion matrix serta precision,

recall, dan F1-score per kelas beserta rata-rata makro pada data uji internal, kemudian memvalidasi generalisasi pada 105 data primer. Interpretabilitas dianalisis menggunakan SHAP pada tingkat global maupun individual [10].

4. Implementasi Sistem

Sistem diimplementasikan dengan arsitektur dua layanan, yaitu aplikasi web berbasis Laravel dan MySQL serta model yang disajikan sebagai layanan inferensi Python yang diakses melalui REST API. Pemisahan ini memungkinkan proses skrining tetap berjalan dengan aturan Permenkes apabila layanan model tidak tersedia (graceful fallback). Pengembangan perangkat lunak mengikuti model Waterfall [21].

4. HASIL DAN PEMBAHASAN

Pada data uji internal sebanyak 1.500 baris, model mencapai akurasi 0,8827 dan F1-macro 0,8580. Metrik per kelas disajikan pada Tabel 1. Kelas Sesuai dikenali dengan precision tertinggi (0,96), sedangkan kelas Meragukan memiliki precision terendah (0,69) namun recall tinggi (0,85), yang menunjukkan model cenderung lebih banyak menggolongkan kasus ke kategori Meragukan.

Tabel 1. Metrik evaluasi per kelas pada data uji

Kelas	Precisi on	Recall	F1-score	Support
Sesuai	0,96	0,91	0,94	867
Meragukan	0,69	0,85	0,76	326
Penyimpangan	0,93	0,83	0,88	307
Rata-rata makro	0,86	0,86	0,86	1500

Distribusi kesalahan klasifikasi disajikan pada Tabel 2. Temuan penting secara klinis adalah bahwa dari 307 kasus Penyimpangan aktual, tidak ada satu pun yang salah diklasifikasikan sebagai Sesuai; seluruh kesalahan pada kelas ini jatuh ke kategori Meragukan yang tetap mengarahkan anak untuk pemeriksaan ulang. Secara keseluruhan tidak terdapat kesalahan yang melompati dua kelas sekaligus—seluruh kesalahan hanya bergeser satu tingkat ke kelas yang bersebelahan—sehingga model meminimalkan kesalahan yang paling berisiko. Kelas Meragukan paling sulit dipisahkan karena menempati rentang skor sempit (7–8) yang berbatasan dengan kedua kelas lain.

Tabel 2. Confusion matrix pada data uji (baris = aktual, kolom = prediksi)

Aktual \ Prediksi	Sesuai	Meragukan	Penyimpangan
Sesuai	793	74	0
Meragukan	31	276	19
Penyimpangan	0	52	255

Pada validasi eksternal menggunakan 105 data primer skrining kelompok usia enam bulan dari Puskesmas Ciampel, prediksi model konsisten dengan penilaian tenaga kesehatan pada seluruh data (akurasi 100%). Angka ini perlu dimaknai hati-hati karena distribusi data primer sangat timpang, yaitu 103 data berkategori Sesuai dan 2 Meragukan tanpa kasus Penyimpangan, dengan seluruh skor berada pada rentang 8–10. Validasi awal ini mengonfirmasi kemampuan generalisasi pada kasus umum, tetapi belum dapat menilai performa pada kelas Meragukan dan Penyimpangan. Dari sisi pola jawaban, butir yang paling sering dijawab Tidak adalah kemampuan berbalik badan (aspek gerak kasar) pada 22 dari 105 anak, sejalan dengan literatur bahwa keterampilan tersebut umumnya masih berkembang pada usia enam bulan. Mengingat ukuran dan cakupan usia data primer yang terbatas, hasil ini diposisikan sebagai indikasi awal generalisasi, bukan estimasi performa yang menyeluruh.

Analisis independensi dilakukan untuk menguji apakah model sekadar mereplikasi aturan ambang skor. Pada data uji sintesis, kesepakatan prediksi model terhadap label aturan ambang sebesar 89,60%, kesepakatan model terhadap label acuan 88,27%, dan kesepakatan aturan ambang terhadap label acuan 84,87%. Selisih 3,40 poin persentase mengindikasikan model menangkap sebagian pola lintas-aspek yang tidak tertangkap penjumlahan skor sederhana, dengan seluruh ketidaksepakatan model hanya bergeser satu tingkat. Karena selisih ini diperoleh pada data sintesis dan pada data primer model maupun aturan ambang sama-sama selaras dengan penilaian tenaga kesehatan, keunggulan model terhadap aturan ambang diposisikan sebagai indikasi awal yang masih memerlukan konfirmasi pada data nyata yang lebih beragam.

Analisis feature importance menunjukkan bahwa jawaban pada aspek bicara-bahasa memberikan kontribusi terbesar terhadap prediksi model (importance 0,183, jauh

melampaui fitur lain pada rentang 0,063–0,087), konsisten dengan posisi literatur tumbuh kembang yang menempatkan keterlambatan bicara sebagai penanda dini. Tingginya kontribusi ini juga dapat dijelaskan dari desain data: mekanisme label drift secara eksplisit menjadikan kelemahan aspek bicara-bahasa sebagai salah satu kondisi penyesuaian label, sehingga konsistensi antara desain pelabelan, feature importance, dan literatur menjadi verifikasi internal bahwa model mempelajari pola yang dirancang. Pada tingkat individual, SHAP menguraikan tiap prediksi dari sebuah nilai dasar dengan menambahkan kontribusi tiap fitur, dan sifat aditivitasnya memastikan penjumlahan nilai dasar dan seluruh kontribusi sama dengan probabilitas kelas prediksi. Penjelasan ini merupakan penjelasan atas keputusan model, bukan pernyataan sebab-akibat klinis.

Temuan penelitian ini sejalan dengan literatur pada beberapa titik. Performa Random Forest yang memadai pada data tabular skrining konsisten dengan studi komparatif yang menempatkan Random Forest sebagai algoritma terbaik pada klasifikasi tumbuh kembang balita [7], meskipun angka metrik antar-penelitian tidak dapat diperbandingkan langsung karena berbeda instrumen dan sumber data. Dibandingkan penerapan explainable AI pada prediksi stunting tanpa integrasi sistem [8], penelitian ini melangkah lebih jauh pada sisi deployment dengan menyajikan penjelasan SHAP langsung di dalam sistem yang digunakan tenaga kesehatan.

Sistem yang dibangun menyajikan tiga lapisan keluaran yang secara langsung menjawab keterbatasan pelaksanaan skrining manual. Lapisan per-skrining individu menampilkan hasil klasifikasi beserta penjelasan SHAP, menjawab keterbatasan interpretasi ambang skor yang tidak menyertakan dasar keputusan. Lapisan longitudinal per anak menjawab ketiadaan pencatatan elektronik untuk memantau trajektori perkembangan. Lapisan agregat tingkat fasilitas menjawab ketiadaan agregasi otomatis antar-Posyandu. Ketiganya diposisikan sebagai augmentasi terhadap KPSP tanpa mengubah peran instrumen maupun tenaga kesehatan.

5. KESIMPULAN

1. Sistem berbasis web yang mengaugmentasi interpretasi KPSP menggunakan Random Forest dan SHAP berhasil dikembangkan, dengan model mencapai akurasi 0,8827 dan F1-macro 0,8580 pada data uji internal.

2. Seluruh kesalahan klasifikasi hanya bergeser satu tingkat ke kelas bersebelahan dan tidak ada kasus Penyimpangan yang salah dikategorikan sebagai Sesuai, sehingga model meminimalkan kesalahan yang paling berisiko secara klinis.

3. Analisis SHAP menunjukkan kontribusi terbesar pada aspek bicara-bahasa yang konsisten dengan literatur dan desain pelabelan, serta memungkinkan setiap prediksi ditelusuri sebagai penjelasan keputusan model.

4. Sistem menyajikan tiga lapisan keluaran—per-individu, longitudinal, dan agregat—yang menjawab keterbatasan skrining manual, dengan tetap memposisikan sistem sebagai alat bantu, bukan pengganti instrumen KPSP maupun tenaga kesehatan.

Keterbatasan penelitian terletak pada data uji primer yang masih terbatas pada satu kelompok usia dengan distribusi yang sangat didominasi kelas Sesuai, serta penggunaan data sintesis untuk pelatihan. Pengembangan selanjutnya diarahkan pada pengumpulan data primer lintas kelompok umur dan pengujian sistem pada lingkungan operasional.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada Puskesmas Ciampel Kabupaten Karawang atas dukungan data dan pelaksanaan penelitian, serta kepada Program Studi Informatika Fakultas Ilmu Komputer Universitas Singaperbangsa Karawang atas dukungan akademik selama penelitian berlangsung.

DAFTAR PUSTAKA

- [1] World Health Organization, World Health Statistics 2018: Monitoring Health for the SDGs. Geneva: World Health Organization, 2018.
- [2] Kementerian Kesehatan Republik Indonesia, Profil Kesehatan Indonesia Tahun 2020. Jakarta: Kementerian Kesehatan Republik Indonesia, 2021.
- [3] Badan Penelitian dan Pengembangan Kesehatan, Laporan Nasional RISKESDAS 2018. Jakarta: Lembaga Penerbit Balitbangkes, 2019.

- [4] World Health Organization and United Nations Children's Fund, *Global Report on Children with Developmental Disabilities*. Geneva: World Health Organization, 2023.
- [5] Kementerian Kesehatan Republik Indonesia, *Peraturan Menteri Kesehatan Republik Indonesia Nomor 66 Tahun 2014 tentang Pemantauan Pertumbuhan, Perkembangan, dan Gangguan Tumbuh Kembang Anak*. Jakarta: Kementerian Kesehatan Republik Indonesia, 2014.
- [6] N. Apriningrum, Carudin, and M. A. Rahayu, "Rancang bangun aplikasi KPSP berbasis Android bagi anak balita sampai pra sekolah di Kabupaten Karawang," *JUSTIN (Jurnal Sistem dan Teknologi Informasi)*, vol. 6, no. 4, pp. 200–204, 2018, doi: 10.26418/justin.v6i4.27385.
- [7] Z. I. Bimawan, T. Astuti, and P. Arsi, "Comparison of random forest, K-nearest neighbor, decision tree, and XGBoost algorithms for detecting stunting in toddlers," *Jurnal Teknik Informatika (JUTIF)*, vol. 5, no. 6, pp. 1599–1607, 2024, doi: 10.52436/1.jutif.2024.5.6.2629.
- [8] A. M. Majid and I. Nawangsih, "Penerapan explainable AI dan model stacking untuk mengidentifikasi faktor risiko stunting balita," *Progresif: Jurnal Ilmiah Komputer*, vol. 22, no. 1, 2026, doi: 10.35889/progresif.v22i1.3549.
- [9] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001, doi: 10.1023/A:1010933404324.
- [10] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Advances in Neural Information Processing Systems 30 (NIPS 2017)*, 2017, pp. 4765–4774.
- [11] R. Wirth and J. Hipp, "CRISP-DM: Towards a standard process model for data mining," in *Proc. 4th Int. Conf. Practical Applications of Knowledge Discovery and Data Mining*, 2000, pp. 29–39.
- [12] R. J. Chen, M. Y. Lu, T. Y. Chen, D. F. K. Williamson, and F. Mahmood, "Synthetic data in machine learning for medicine and healthcare," *Nature Biomedical Engineering*, vol. 5, no. 6, pp. 493–497, 2021, doi: 10.1038/s41551-021-00751-8.
- [13] T. van der Ploeg, P. C. Austin, and E. W. Steyerberg, "Modern modelling techniques are data hungry: A simulation study for predicting dichotomous endpoints," *BMC Medical Research Methodology*, vol. 14, art. 137, 2014, doi: 10.1186/1471-2288-14-137.
- [14] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Transactions on Knowledge and Data Engineering*, vol. 21, no. 9, pp. 1263–1284, 2009, doi: 10.1109/TKDE.2008.239.
- [15] R. Kohavi, "A study of cross-validation and bootstrap for accuracy estimation and model selection," in *Proc. 14th Int. Joint Conf. Artificial Intelligence (IJCAI)*, 1995, pp. 1137–1143.
- [16] P. Probst, M. N. Wright, and A.-L. Boulesteix, "Hyperparameters and tuning strategies for random forest," *WIREs Data Mining and Knowledge Discovery*, vol. 9, no. 3, art. e1301, 2019, doi: 10.1002/widm.1301.
- [17] T. M. Oshiro, P. S. Perez, and J. A. Baranauskas, "How many trees in a random forest?," in *Machine Learning and Data Mining in Pattern Recognition (LNCS 7376)*. Berlin: Springer, 2012, pp. 154–168.
- [18] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, 2nd ed. New York, NY: Springer, 2009.
- [19] M. H. P. Andriano, E. Purwanto, and D. Hartanti, "Sistem untuk monitoring data tumbuh kembang anak menggunakan klasifikasi decision tree di PAUD IT Nur Hidayah," in *Prosiding Seminar Nasional Teknologi Informasi dan Bisnis (SENATIB) 2025*, Surakarta, 2025, pp. 290–297.
- [20] I. M. Rahayu, A. Yusuf, and M. Ridwan, "Prediksi kesiapan sekolah menggunakan machine learning berbasis kombinasi Adam dan Nesterov Momentum," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 9, no. 6, pp. 1273–1280, 2022.
- [21] I. Sommerville, *Software Engineering*, 9th ed. Boston, MA: Addison-Wesley, 2011.
- [22] D. Duei Putri, G. F. Nama, and W. E. Sulistiono, "Analisis sentimen kinerja Dewan Perwakilan Rakyat (DPR) pada Twitter menggunakan metode Naive Bayes Classifier," *JITET (Jurnal Informatika dan Teknik Elektro Terapan)*, vol. 10, no. 1, pp. 34–40, 2022, doi: 10.23960/jitet.v10i1.2262.