

BAB II

LANDASAN TEORI

2.1 Penelitian Terdahulu

Penelitian terdahulu digunakan untuk memahami perkembangan penelitian terkait prediksi stunting menggunakan machine learning serta mengidentifikasi celah penelitian yang belum banyak dikaji. Berdasarkan pedoman program studi, penelitian terdahulu yang digunakan minimal 10 jurnal dalam 5 tahun terakhir.

Tabel 2.1 Jurnal Penelitian Terdahulu

Referensi ke-1	
Judul	Analisa Optimasi Grid Search pada Algoritma Random Forest dan Decision Tree untuk Klasifikasi Stunting
Nama Penulis	Rahmayani & Budiman [28]
Sumber Jurnal	Building of Informatics, Technology and Science (BITS), Vol. 6, No. 3
Tahun	2024
Permasalahan	Penelitian ini bertujuan mengatasi belum optimalnya performa algoritma Decision Tree dan Random Forest dalam klasifikasi stunting. Penelitian berfokus pada peningkatan akurasi model melalui optimasi hyperparameter menggunakan metode Grid Search sehingga dapat diperoleh model klasifikasi yang lebih optimal.
Framework / Model	Data Mining, Data Preprocessing (Data Selection, Data Cleaning, Mapping, Class Balancing menggunakan SMOTE, Standard Scaler), Decision Tree, Random Forest, Grid Search, Confusion Matrix.
Temuan	Hasil penelitian menunjukkan bahwa optimasi menggunakan Grid Search mampu meningkatkan performa kedua algoritma. Akurasi Decision Tree meningkat dari 70,2% menjadi 82,8%, sedangkan Random Forest meningkat dari 77,9% menjadi 84,1%. Random Forest menghasilkan performa terbaik setelah proses optimasi.
Pembahasan	Penelitian menunjukkan bahwa optimasi hyperparameter melalui Grid Search berpengaruh terhadap peningkatan performa model klasifikasi stunting. Namun penelitian menggunakan dataset publik dari Kaggle dan berfokus pada peningkatan akurasi model melalui optimasi algoritma tanpa mengimplementasikan model dalam

	prototype berbasis website maupun melakukan validasi pengguna.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena sama-sama menerapkan machine learning untuk klasifikasi status stunting menggunakan data antropometri. Perbedaanannya, penelitian ini menggunakan data operasional balita dari wilayah kerja Puskesmas Pagedangan, membandingkan algoritma K-Nearest Neighbor dan Support Vector Machine, serta mengimplementasikan model terbaik ke dalam prototype berbasis website/Flask sebagai media simulasi alat bantu klasifikasi awal status stunting.
Referensi ke-2	
Judul	Designing a Stunting Prediction Model Using Machine Learning to Support SDGs Achievement in Indonesia
Nama Penulis	Sinaga et al. [4]
Sumber Jurnal	SinkrOn: Jurnal dan Penelitian Teknik Informatika, Volume 9, Issue 4
Tahun	2025
Permasalahan	Penelitian bertujuan mengatasi keterbatasan metode pemantauan stunting yang masih bersifat reaktif sehingga diperlukan model prediksi berbasis machine learning untuk mendukung deteksi dini stunting serta pencapaian Sustainable Development Goals (SDGs).
Framework / Model	Data Collection, Data Preprocessing, Feature Selection, Model Development, Evaluation and Interpretation menggunakan algoritma K-Nearest Neighbor (KNN), Naïve Bayes, Decision Tree, dan Random Forest, dengan 10-Fold Cross Validation serta McNemar's Test sebagai evaluasi statistik.
Temuan	Decision Tree memperoleh performa terbaik dengan accuracy 88,58%, diikuti Random Forest 88,32%, KNN 84,70%, dan Naïve Bayes 72%. Hasil uji McNemar menunjukkan Decision Tree dan Random Forest tidak memiliki perbedaan performa yang signifikan.
Pembahasan	Penelitian menunjukkan bahwa algoritma berbasis pohon memberikan performa lebih baik dibandingkan KNN dan Naïve Bayes pada data stunting yang memiliki distribusi kelas tidak seimbang. Evaluasi dilakukan menggunakan accuracy, precision, recall, F1-score, ROC Area, dan McNemar's Test untuk memastikan keandalan model.
Relevansi dan Gap terhadap Penelitian Ini	Relevan karena sama-sama menerapkan machine learning untuk klasifikasi atau prediksi stunting berdasarkan data antropometri. Namun penelitian menggunakan dataset regional Sumatera Utara dengan empat algoritma klasifikasi, sedangkan penelitian ini menggunakan data operasional Puskesmas Pagedangan, membandingkan

	algoritma K-Nearest Neighbor (KNN) dan Support Vector Machine (SVM), serta mengimplementasikan model terbaik ke dalam prototype berbasis website/Flask sebagai media simulasi alat bantu klasifikasi awal status stunting
Referensi ke-3	
Judul	Benchmarking machine learning algorithm for stunting risk prediction in Indonesia
Nama Penulis	Novalina et al. [29]
Sumber Jurnal	<i>Bulletin of Electrical Engineering and Informatics (BEEI)</i> , Vol. 14, No. 3
Tahun	2025
Permasalahan	Penelitian bertujuan mengembangkan model prediksi risiko stunting di Indonesia menggunakan pendekatan machine learning. Permasalahan utama yang diangkat adalah tingginya prevalensi stunting serta perlunya identifikasi faktor-faktor dominan yang memengaruhi stunting melalui pendekatan data science. Penelitian juga berupaya mengatasi permasalahan ketidakseimbangan kelas (class imbalance) pada dataset agar menghasilkan model prediksi yang lebih akurat.
Framework / Model	Penelitian menggunakan framework CRISP-DM yang terdiri dari Business Understanding, Data Understanding, Data Preparation, Modeling, Evaluation, dan Deployment. Tahap preprocessing meliputi feature selection, Random Undersampling, dan SMOTE. Model yang dibandingkan adalah Logistic Regression, K-Nearest Neighbor (KNN), Support Vector Classifier (SVC), dan Decision Tree dengan evaluasi menggunakan Accuracy, Precision, Recall, F1-score, serta implementasi GUI sebagai tahap deployment.
Temuan	Logistic Regression menghasilkan performa terbaik setelah menggunakan SMOTE dengan Accuracy 96,17%, Precision 96,44%, Recall 96,17%, dan F1-score 96,25%. Penggunaan SMOTE terbukti memberikan performa lebih baik dibandingkan Random Undersampling pada seluruh model yang diuji.
Pembahasan	Penelitian menunjukkan bahwa penerapan CRISP-DM, penanganan data tidak seimbang menggunakan SMOTE, serta proses feature selection mampu meningkatkan performa model klasifikasi. Selain membandingkan beberapa algoritma machine learning, penelitian juga mengimplementasikan model ke dalam GUI sebagai media prediksi risiko stunting berdasarkan variabel antropometri dan faktor pendukung lainnya.

Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini sangat relevan karena sama-sama membahas klasifikasi/prediksi stunting menggunakan machine learning, menerapkan framework CRISP-DM, menggunakan teknik SMOTE, serta melakukan deployment prototype. Namun penelitian ini menggunakan data Indonesia Family Life Survey (IFLS) dengan fokus pada prediksi risiko stunting menggunakan empat algoritma. Sementara itu, penelitian yang dilakukan menggunakan data antropometri operasional Puskesmas Pagedangan, membandingkan K-Nearest Neighbor (KNN) dan Support Vector Machine (SVM) secara khusus, serta memanfaatkan prototype berbasis Flask sebagai alat bantu klasifikasi status stunting bagi tenaga kesehatan
Referensi ke-4	
Judul	Comparison of Random Forest, K-Nearest Neighbor, Decision Tree, and XGBoost Algorithms for Detecting Stunting in Toddlers
Nama Penulis	Bimawan et al. [30]
Sumber Jurnal	Jurnal Teknik Informatika (JUTIF), Vol. 5, No. 6
Tahun	2024
Permasalahan	Penelitian membahas tingginya angka stunting di Indonesia dan pentingnya deteksi dini menggunakan machine learning. Penelitian menilai belum adanya perbandingan komprehensif terhadap beberapa algoritma klasifikasi untuk mendeteksi stunting pada balita sehingga diperlukan evaluasi performa masing-masing algoritma.
Framework / Model	Tahapan penelitian meliputi pengumpulan data, eksplorasi data, pra-pemrosesan data, ekstraksi fitur, klasifikasi model, dan evaluasi model. Algoritma yang dibandingkan adalah Random Forest, K-Nearest Neighbor (KNN), Decision Tree, dan XGBoost. Evaluasi menggunakan Accuracy, Precision, Recall, dan F1-score. Dataset dibagi menjadi 80% data training dan 20% data testing.
Temuan	Random Forest memberikan performa terbaik dengan Accuracy 99,9132%, Recall 99,9132%, dan Macro F1-score 99,8906%. KNN dan Decision Tree juga menunjukkan performa yang sangat tinggi, sedangkan XGBoost memiliki performa sedikit lebih rendah dibandingkan ketiga algoritma lainnya.
Pembahasan	Penelitian menunjukkan seluruh algoritma memiliki performa klasifikasi yang sangat baik, namun Random Forest menjadi model paling konsisten berdasarkan seluruh metrik evaluasi. Penelitian juga menyajikan confusion matrix setiap algoritma sebagai dasar analisis hasil klasifikasi status gizi balita.

Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini relevan karena sama-sama membandingkan algoritma machine learning untuk klasifikasi status stunting dan menggunakan KNN sebagai salah satu model. Namun penelitian ini menggunakan Random Forest, Decision Tree, KNN, dan XGBoost dengan dataset dari Kaggle, sedangkan penelitian yang dilakukan menggunakan data antropometri Puskesmas Pagedangan, membandingkan KNN dan SVM, menerapkan framework CRISP-DM, serta melakukan optimasi parameter menggunakan GridSearchCV dan implementasi prototype berbasis Flask.
Referensi ke-5	
Judul	Deteksi Dini Stunting Pada Anak Berdasarkan Indikator Antropometri dengan Menggunakan Algoritma Machine Learning
Nama Penulis	Ratnasari et al. [31]
Sumber Jurnal	Jurnal Algoritma, Vol. 21 No. 2
Tahun	2024
Permasalahan	Penelitian membahas pentingnya deteksi dini stunting menggunakan data antropometri karena metode konvensional masih bergantung pada pemeriksaan manual yang membutuhkan waktu lebih lama. Penelitian bertujuan membangun model machine learning yang mampu meningkatkan efektivitas skrining stunting pada anak.
Framework / Model	Penelitian menggunakan tahapan pra-pemrosesan data, normalisasi, pemodelan machine learning, dan evaluasi model. Algoritma yang dibandingkan meliputi Random Forest, K-Nearest Neighbor (KNN), Support Vector Machine (SVM), dan Naive Bayes. Evaluasi dilakukan menggunakan Accuracy, Confusion Matrix, dan ROC Curve dengan 10-fold Cross Validation.
Temuan	Random Forest memperoleh performa terbaik dengan akurasi 92,70%, diikuti KNN sebesar 91,40%, sedangkan SVM dan Naive Bayes masing-masing memperoleh akurasi 86,60%. Random Forest direkomendasikan sebagai model terbaik untuk deteksi dini stunting.
Pembahasan	Random Forest memiliki kemampuan yang lebih baik dalam menangani hubungan nonlinier antar fitur antropometri sehingga menghasilkan performa tertinggi. Penelitian juga menunjukkan bahwa tinggi badan dan usia merupakan fitur yang paling berpengaruh dalam prediksi stunting serta membahas tantangan implementasi model pada fasilitas kesehatan.
Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini sangat relevan karena sama-sama menggunakan indikator antropometri dan membandingkan algoritma machine learning termasuk KNN dan SVM

	untuk klasifikasi stunting. Namun penelitian ini menggunakan empat algoritma (Random Forest, KNN, SVM, dan Naive Bayes) dengan 220 data balita dari Puskesmas Wates Kabupaten Pringsewu, sedangkan penelitian yang dilakukan berfokus pada perbandingan KNN dan SVM menggunakan data antropometri Puskesmas Pagedangan, menerapkan framework CRISP-DM, serta melakukan optimasi parameter menggunakan GridSearchCV sebelum membandingkan performa kedua model.
Referensi ke-6	
Judul	Classification Analysis of Stunting Data on Toddlers in Bekasi City Using Logistic Regression and Random Forest
Nama Penulis	Nadhira & Gunawan [32]
Sumber Jurnal	2025 International Conference on Data Science and Its Applications (ICoDSA)
Tahun	2025
Permasalahan	Penelitian bertujuan mengembangkan model klasifikasi status stunting pada balita di Kota Bekasi menggunakan algoritma machine learning. Permasalahan yang diangkat adalah tingginya prevalensi stunting di Indonesia sehingga diperlukan model prediksi yang lebih akurat untuk mendukung pemantauan pertumbuhan dan perkembangan balita.
Framework / Model	Tahapan penelitian meliputi Literature Review, Data Collection, Data Cleaning, Feature Selection, Data Visualization, Data Balancing, Logistic Regression dan Random Forest Modeling, Accuracy Comparison, serta Conclusion. Evaluasi model menggunakan 5-Fold Cross Validation dan Confusion Matrix.
Temuan	Hasil penelitian menunjukkan bahwa Random Forest memperoleh performa terbaik dengan Accuracy sebesar 88,07% dan Macro-average F1-score sebesar 87,32%, sedangkan Logistic Regression memperoleh Accuracy sebesar 85,48% dan Macro-average F1-score sebesar 84,73%.
Pembahasan	Penelitian menunjukkan bahwa Random Forest memberikan performa klasifikasi yang lebih baik dibandingkan Logistic Regression dalam mendeteksi status stunting balita. Selain itu, proses data balancing dan validasi menggunakan 5-Fold Cross Validation membantu meningkatkan reliabilitas model dalam menangani distribusi data yang tidak seimbang.
Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini relevan karena sama-sama menerapkan machine learning untuk klasifikasi status stunting

	berdasarkan data balita. Namun penelitian ini hanya membandingkan Logistic Regression dan Random Forest, sedangkan penelitian yang dilakukan menggunakan K-Nearest Neighbor (KNN) dan Support Vector Machine (SVM) dengan data antropometri operasional dari Puskesmas Pagedangan. Selain itu, penelitian ini juga mengimplementasikan model terbaik ke dalam prototype berbasis website/Flask sebagai media simulasi alat bantu klasifikasi awal status stunting.
Referensi ke-7	
Judul	Posyandu Application for Monitoring Children Under-Five: A 3-Year Data Quality Map in Indonesia
Nama Penulis	Rinawan et al. [12]
Sumber Jurnal	SPRS International Journal of Geo-Information (MDPI), Vol. 11, Issue 7
Tahun	2022
Permasalahan	Penelitian membahas belum optimalnya kualitas data Posyandu di Indonesia sehingga diperlukan evaluasi terhadap aspek kelengkapan (completeness), akurasi (accuracy), dan konsistensi (consistency) data selama tiga tahun. Kualitas data yang baik diperlukan agar dapat digunakan sebagai dasar analisis dan pengambilan keputusan kesehatan masyarakat
Framework / Model	Menggunakan WHO Data Quality Review (DQR) Framework dengan tiga dimensi, yaitu completeness, accuracy, dan internal consistency. Analisis dilakukan menggunakan WHO Anthropometry Z-Score, Cronbach's Alpha, STATA, dan QGIS untuk memetakan kualitas data Posyandu di Indonesia.
Temuan	Data Posyandu pada wilayah pilot project memiliki kualitas yang lebih baik dibanding wilayah lainnya. Pandemi menyebabkan penurunan kelengkapan, akurasi, dan konsistensi data. Penelitian juga menegaskan bahwa peningkatan kualitas data sangat penting sebelum data dimanfaatkan untuk analisis lanjutan maupun machine learning.
Pembahasan	Penelitian menjelaskan bahwa kualitas data merupakan fondasi utama dalam analisis kesehatan. Faktor seperti pelatihan kader Posyandu, proses validasi data, pengawasan bidan, serta integrasi aplikasi iPosyandu dengan sistem pemerintah berpengaruh terhadap kualitas data. Pada bagian diskusi juga dijelaskan bahwa machine learning memerlukan data yang telah melalui proses preprocessing dan validasi agar menghasilkan model yang akurat.

Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini relevan karena menggunakan data antropometri Posyandu serta menekankan pentingnya kualitas data sebelum dilakukan analisis machine learning. Namun penelitian tersebut tidak membangun model klasifikasi stunting maupun membandingkan algoritma machine learning. Penelitian yang dilakukan menggunakan K-Nearest Neighbor (KNN) dan Support Vector Machine (SVM) untuk mengklasifikasikan status stunting berdasarkan data antropometri balita, sehingga memberikan kontribusi pada tahap analisis prediktif, bukan hanya evaluasi kualitas data
Referensi ke-8	
Judul	Exploring practical issues in children's anthropometric measurements: A qualitative descriptive study involving Indonesian health professionals and community health workers.
Nama Penulis	Wanda et al. [13]
Sumber Jurnal	Belitung Nursing Journal, Vol. 11, No. 5
Tahun	2025
Permasalahan	Penelitian mengidentifikasi berbagai kendala dalam proses pengukuran antropometri balita di Posyandu yang dapat menyebabkan data pertumbuhan menjadi kurang akurat. Permasalahan tersebut meliputi keterampilan kader yang tidak merata, kesalahan penggunaan alat ukur, kondisi alat yang kurang terkalibrasi, perilaku anak dan orang tua saat pengukuran, serta keterbatasan fasilitas Posyandu. Kondisi tersebut berpotensi menghasilkan data yang kurang valid untuk menentukan status gizi maupun prevalensi stunting.
Framework / Model	Menggunakan pendekatan kualitatif deskriptif (Qualitative Descriptive Study) dengan Focus Group Discussion (FGD) terhadap 10 tenaga kesehatan dan 8 kader Posyandu. Analisis data dilakukan menggunakan Thematic Analysis berdasarkan metode Braun & Clarke.
Temuan	Penelitian menghasilkan empat tema utama, yaitu barriers to measurement accuracy, varied skills in measurement, mothers' behavior influenced by children's reactions, serta strategies for dealing with traumatized children. Faktor-faktor tersebut terbukti memengaruhi akurasi pengukuran antropometri sehingga berpengaruh terhadap kualitas data pertumbuhan anak.
Pembahasan	Penelitian menegaskan bahwa kualitas data antropometri dipengaruhi oleh kalibrasi alat, kompetensi kader Posyandu, pelatihan berkala, kondisi lingkungan pengukuran, serta komunikasi dengan orang tua. Penulis juga merekomendasikan penggunaan alat ukur yang tervalidasi, pelatihan rutin bagi kader, serta

	penyederhanaan prosedur pengukuran agar data yang dihasilkan lebih akurat dan dapat digunakan sebagai dasar pengambilan keputusan kesehatan.
Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini relevan karena menggunakan konteks Posyandu di Indonesia dan menyoroti pentingnya kualitas data antropometri sebagai dasar penentuan status gizi dan stunting. Namun, penelitian tersebut tidak melakukan klasifikasi status stunting menggunakan machine learning maupun membandingkan performa algoritma klasifikasi. Penelitian yang dilakukan menggunakan K-Nearest Neighbor (KNN) dan Support Vector Machine (SVM) untuk mengklasifikasikan status stunting berdasarkan data antropometri sehingga memberikan kontribusi pada analisis prediktif, bukan hanya evaluasi proses pengukuran.
Referensi ke-9	
Judul	The Impact of the Inaccuracy Measurement of Anthropometry by Posyandu Cadres on the Classification of Stunting of Children Under Five Years Old
Nama Penulis	Suyatno et al. [14]
Sumber Jurnal	Annals of Tropical Medicine and Public Health, Vol. 24, No. 1
Tahun	2021
Permasalahan	Ketepatan pengukuran antropometri oleh kader Posyandu menjadi faktor penting dalam penentuan status stunting. Kesalahan pengukuran tinggi badan maupun panjang badan berpotensi menyebabkan klasifikasi status stunting menjadi tidak akurat.
Framework / Model	Evaluasi akurasi pengukuran antropometri oleh kader Posyandu melalui perbandingan hasil pengukuran dengan pengukuran standar tenaga kesehatan.
Temuan	Penelitian menunjukkan bahwa ketidakakuratan pengukuran antropometri oleh kader Posyandu dapat menyebabkan kesalahan klasifikasi status stunting pada balita. Oleh karena itu, kompetensi kader dan standar prosedur pengukuran menjadi faktor penting dalam menjaga kualitas data.
Pembahasan	Penelitian menekankan pentingnya pelatihan kader, penggunaan prosedur pengukuran yang sesuai standar WHO, serta pengawasan terhadap proses pengukuran untuk meminimalkan kesalahan klasifikasi status stunting. Fokus penelitian berada pada kualitas proses pengukuran, bukan pada pengembangan sistem digital ataupun model machine learning.

Relevansi dan Gap terhadap Penelitian Ini	Penelitian ini memperkuat dasar bahwa kualitas data antropometri sangat memengaruhi hasil klasifikasi status stunting. Namun penelitian tersebut belum memanfaatkan machine learning untuk melakukan klasifikasi maupun sebagai alat bantu peninjauan data. Penelitian yang dilakukan melengkapi gap tersebut dengan membandingkan algoritma K-Nearest Neighbor (KNN) dan Support Vector Machine (SVM) untuk mengklasifikasikan status stunting berdasarkan data antropometri.
Referensi ke-10	
Judul	A survey on dataset quality in machine learning
Nama Penulis	Gong et al. [18]
Sumber Jurnal	<i>Information and Software Technology</i> , Vol. 162, Elsevier
Tahun	2023
Permasalahan	Banyak model machine learning memiliki performa rendah karena kualitas dataset yang buruk, seperti missing value, duplicate data, inconsistent data, incorrect data, noise, dan ketidakseimbangan data. Penelitian mengkaji bagaimana kualitas dataset memengaruhi performa model machine learning.
Framework / Model	Systematic Literature Review (Survey Paper) mengenai Dataset Quality, Data Quality Evaluation, Data Preparation, dan Machine Learning
Temuan	Kualitas dataset memiliki pengaruh yang sangat besar terhadap akurasi, stabilitas, dan kemampuan generalisasi model machine learning. Tahapan preprocessing seperti data cleaning, handling missing value, duplicate removal, data consistency checking, dan quality assessment menjadi faktor penting sebelum proses pemodelan dilakukan
Pembahasan	Penelitian menjelaskan berbagai dimensi kualitas data, metode evaluasi kualitas dataset, serta teknik preprocessing yang dapat meningkatkan performa model machine learning. Paper ini menekankan bahwa model yang baik tidak hanya ditentukan oleh algoritma, tetapi juga oleh kualitas dataset yang digunakan
Relevansi dan Gap terhadap Penelitian Ini	Relevansi: Mendukung pentingnya tahap Data Preparation pada framework CRISP-DM yang digunakan dalam penelitian ini sebelum membangun model KNN dan SVM menggunakan data antropometri balita. Gap: Penelitian ini bersifat survei dan tidak menerapkan algoritma klasifikasi maupun studi kasus pada data stunting. Penelitian yang dilakukan menggunakan dataset antropometri balita dari Puskesmas serta membandingkan performa algoritma KNN dan SVM setelah melalui proses preprocessing

Penelitian terdahulu digunakan sebagai dasar untuk memahami perkembangan penerapan machine learning dalam klasifikasi status stunting, kualitas data antropometri, serta pemanfaatan data kesehatan sebagai pendukung pengambilan keputusan. Berdasarkan hasil kajian terhadap sepuluh penelitian terdahulu, diketahui bahwa berbagai algoritma telah digunakan untuk klasifikasi maupun prediksi status stunting, seperti Decision Tree, Random Forest, K-Nearest Neighbor (KNN), Support Vector Machine (SVM), Logistic Regression, Naïve Bayes, dan XGBoost. Hasil pemetaan tersebut juga menunjukkan bahwa performa model dipengaruhi oleh karakteristik dataset, tahapan preprocessing, distribusi kelas, metode penanganan ketidakseimbangan data, serta metode evaluasi yang digunakan [4], [28], [29], [30], [31], [32]. Selain itu, penelitian terkait kualitas data Posyandu, pengukuran antropometri, dan kualitas dataset machine learning menegaskan bahwa kesalahan pengukuran, ketidakkonsistenan pencatatan, missing value, data tidak valid, duplikasi, serta ketidakseimbangan kelas dapat memengaruhi kualitas analisis dan hasil klasifikasi status gizi atau status stunting balita [12], [13], [14], [18].

Berdasarkan pemetaan tersebut, gap penelitian ini terletak pada penggunaan data operasional Puskesmas, evaluasi model pada data yang tidak seimbang, keterkaitan antara kualitas data antropometri dan proses pemodelan, serta pemanfaatan prototype sebagai alat bantu klasifikasi awal. Oleh karena itu, kontribusi penelitian ini terletak pada pengembangan model klasifikasi status stunting berbasis data antropometri Puskesmas Pagedangan dengan membandingkan algoritma KNN dan SVM. KNN dipilih karena merepresentasikan pendekatan klasifikasi berbasis kedekatan jarak antar data, sedangkan SVM merepresentasikan pendekatan klasifikasi berbasis pembentukan hyperplane dengan margin optimal. Pengujian kedua algoritma pada data Puskesmas Pagedangan yang memiliki distribusi kelas tidak seimbang memberikan kontribusi empiris terhadap pemahaman performa model machine learning pada data layanan kesehatan primer. Selain itu, penelitian ini menambahkan tahap deployment dalam bentuk prototype Flask sebagai simulasi alat bantu klasifikasi awal status stunting yang tetap memerlukan verifikasi lebih lanjut oleh tenaga kesehatan.

2.2 Teori yang Berkaitan

Teori yang berkaitan digunakan sebagai landasan konseptual dalam penelitian ini. Bagian ini membahas konsep stunting, faktor penyebab dan dampaknya, standar pengukuran status stunting, peran Posyandu dalam pemantauan pertumbuhan balita, pencatatan data antropometri, kualitas data kesehatan, serta konsep dasar machine learning dan klasifikasi. Teori-teori tersebut digunakan untuk mendukung pemahaman terhadap permasalahan penelitian dan metode yang digunakan dalam proses klasifikasi status stunting balita.

2.2.1 Stunting

Stunting merupakan kondisi gagal tumbuh pada anak balita akibat kekurangan gizi kronis, infeksi berulang, serta faktor lingkungan yang berlangsung dalam jangka panjang. Kondisi ini ditandai dengan tinggi badan atau panjang badan anak yang lebih rendah dibandingkan standar pertumbuhan anak seusianya. Secara antropometri, anak dikategorikan mengalami stunting apabila nilai *z-score* tinggi badan menurut umur atau panjang badan menurut umur berada di bawah -2 standar deviasi dari median standar pertumbuhan WHO. Anak dengan nilai *z-score* di bawah -3 standar deviasi dikategorikan sebagai *severely stunted* atau sangat pendek. Stunting tidak hanya menggambarkan kondisi tubuh pendek, tetapi juga mencerminkan adanya gangguan pertumbuhan yang terjadi dalam periode panjang. Gangguan tersebut umumnya berkaitan dengan periode 1.000 Hari Pertama Kehidupan, yaitu masa sejak kehamilan hingga anak berusia dua tahun. Periode ini menjadi fase penting karena kekurangan gizi atau gangguan kesehatan yang terjadi pada masa tersebut dapat berdampak terhadap pertumbuhan fisik, perkembangan otak, dan kualitas hidup anak di masa depan [2].

2.2.2 Penyebab Stunting

Penyebab stunting bersifat multifaktorial dan saling berinteraksi, mencakup faktor langsung maupun tidak langsung. Secara umum, penyebab stunting dapat dikelompokkan ke dalam tiga tingkatan.

Penyebab langsung stunting meliputi kekurangan asupan zat gizi dan infeksi berulang. Kekurangan asupan zat gizi seperti energi, protein, zinc, dan zat besi dapat menghambat proses pertumbuhan anak. Secara khusus, asupan protein yang tidak adekuat dapat menghambat sintesis hormon pertumbuhan. Selain itu, infeksi berulang seperti diare, ISPA, dan pneumonia dapat menyebabkan peradangan kronis yang mengganggu penyerapan nutrisi dalam tubuh [33].

Selain faktor langsung, terdapat pula penyebab tidak langsung yang berkontribusi terhadap terjadinya stunting. Faktor ibu memiliki peran yang signifikan, seperti tinggi badan ibu kurang dari 150 cm, tingkat pendidikan ibu yang rendah, serta status gizi ibu selama kehamilan yang kurang baik [33]. Kondisi tersebut terbukti meningkatkan risiko terjadinya stunting pada anak. Selain itu, faktor lain seperti kerawanan pangan keluarga, sanitasi yang buruk, keterbatasan akses terhadap air bersih, serta rendahnya akses terhadap layanan kesehatan juga turut memengaruhi kejadian stunting [34] [35].

Stunting juga dapat dipengaruhi oleh faktor genetik dan faktor lingkungan. Faktor genetik berkaitan dengan potensi pertumbuhan anak yang dapat dipengaruhi oleh karakteristik biologis orang tua, seperti tinggi badan ibu dan ayah. Anak yang lahir dari orang tua dengan postur tubuh pendek dapat memiliki kecenderungan tinggi badan yang lebih rendah. Namun, faktor genetik tidak dapat dijadikan satu-satunya penyebab stunting karena pertumbuhan anak juga sangat dipengaruhi oleh faktor lingkungan [36].

Faktor lingkungan berperan penting dalam menentukan kualitas pertumbuhan anak. Lingkungan yang kurang mendukung, seperti sanitasi yang buruk, keterbatasan akses air bersih, rendahnya kualitas asupan makanan, infeksi berulang, keterbatasan akses layanan kesehatan, serta kondisi sosial ekonomi keluarga dapat meningkatkan risiko terjadinya stunting [37]. Oleh karena itu, stunting perlu dipahami sebagai kondisi yang dipengaruhi oleh interaksi antara faktor biologis, genetik, gizi, kesehatan, dan lingkungan [38].

Pada tingkat yang lebih luas, penyebab mendasar stunting berkaitan dengan kondisi sosial dan ekonomi masyarakat. Faktor-faktor seperti kemiskinan,

rendahnya tingkat pendidikan, ketimpangan gender, serta keterbatasan infrastruktur menjadi determinan utama yang memengaruhi lingkungan tempat anak tumbuh dan berkembang. Kondisi ini secara tidak langsung memengaruhi akses terhadap pangan, layanan kesehatan, dan kualitas lingkungan yang berdampak pada status gizi anak [39].

2.2.3 Dampak Stunting

Stunting dapat menimbulkan dampak jangka pendek maupun jangka panjang. Dalam jangka pendek, stunting dapat meningkatkan risiko kesakitan, menurunkan daya tahan tubuh, serta menghambat perkembangan motorik dan kognitif anak. Anak yang mengalami stunting cenderung memiliki risiko lebih tinggi mengalami keterlambatan perkembangan, sehingga dapat memengaruhi kemampuan belajar sejak usia dini [5], [6].

Dalam jangka panjang, stunting dapat berdampak pada rendahnya prestasi belajar, menurunnya produktivitas kerja, serta meningkatnya risiko penyakit tidak menular pada usia dewasa. Kondisi ini menunjukkan bahwa stunting tidak hanya memengaruhi kesehatan individu, tetapi juga dapat berdampak pada kualitas sumber daya manusia secara lebih luas. Oleh karena itu, deteksi dini dan pemantauan status pertumbuhan balita menjadi bagian penting dalam upaya pencegahan dan penanganan stunting [39].

2.2.4 Standar Pengukuran Stunting

Penilaian status stunting dilakukan menggunakan indeks tinggi badan menurut umur (TB/U) atau panjang badan menurut umur (PB/U). Indeks ini membandingkan tinggi atau panjang badan anak dengan standar pertumbuhan anak berdasarkan umur dan jenis kelamin [40]. Standar yang umum digunakan adalah WHO Child Growth Standards, yang menyediakan tabel dan kurva pertumbuhan anak dalam bentuk z-score [41]. Di Indonesia, penilaian status gizi anak juga mengacu pada Standar Antropometri Anak yang ditetapkan melalui Peraturan Menteri Kesehatan Nomor 2 Tahun 2020 [42].

Pada balita, istilah panjang badan umumnya digunakan untuk anak usia di bawah 24 bulan, sedangkan tinggi badan digunakan untuk anak usia 24 bulan

ke atas. Dalam penelitian ini, dataset menggunakan kolom “Tinggi” dan status TB/U. Oleh karena itu, istilah tinggi atau panjang badan digunakan untuk menjelaskan variabel antropometri secara umum, sedangkan istilah tinggi digunakan ketika merujuk pada nama variabel dalam dataset dan proses pemodelan.

Z-score merupakan nilai standar deviasi yang menunjukkan posisi hasil pengukuran anak terhadap median standar pertumbuhan. Nilai z-score dihitung dengan membandingkan nilai ukur anak terhadap nilai median standar pada kelompok umur dan jenis kelamin yang sama. Rumus z-score dapat dituliskan sebagai berikut:

$$Z - score = \frac{\text{Nilai Ukur Anak} - \text{Nilai Median Standar}}{\text{Nilai Simpang Baku (SD)}} \quad (2.1)$$

Keterangan:

- 1) Nilai Ukur Anak adalah hasil pengukuran tinggi badan atau panjang badan anak.
- 2) Nilai Median Standar adalah nilai median tinggi atau panjang badan berdasarkan standar pertumbuhan anak menurut umur dan jenis kelamin.
- 3) Nilai Simpang Baku (SD) adalah simpangan baku pada standar pertumbuhan yang digunakan.

Semakin rendah nilai z-score TB/U atau PB/U, semakin besar indikasi bahwa tinggi atau panjang badan anak berada di bawah standar pertumbuhan usianya. Anak dengan nilai z-score di bawah -2 SD dikategorikan sebagai stunted atau pendek, sedangkan anak dengan nilai z-score di bawah -3 SD dikategorikan sebagai severely stunted atau sangat pendek. Anak dengan nilai z-score antara -2 SD sampai +3 SD dikategorikan normal, sedangkan anak dengan nilai z-score di atas +3 SD dikategorikan tinggi [43].

Pengukuran stunting menggunakan indikator TB/U (Tinggi Badan menurut Umur) atau PB/U (Panjang Badan menurut Umur) untuk anak di bawah 2 tahun,

yang dibandingkan terhadap standar pertumbuhan WHO 2006. Kategori status gizi berdasarkan z-score TB/U adalah sebagai berikut:

Tabel 2.2 Kategori Status Gizi Berdasarkan Indeks TB/U atau PB/U

Kategori	Z-Score TB/U
Sangat Pendek (<i>Severely Stunted</i>)	$< -3 \text{ SD}$
Pendek (<i>Stunted</i>)	$-3 \text{ SD s/d } -2 \text{ SD}$
Normal	$-2 \text{ SD s/d } +3 \text{ SD}$
Tinggi (<i>Tall</i>)	$> +3 \text{ SD}$

Dalam penelitian ini, standar pengukuran TB/U atau PB/U dan z-score digunakan sebagai dasar teoritis untuk memahami kategori status stunting, yaitu normal, stunted, dan severely stunted. Namun, penelitian ini tidak menghitung ulang nilai z-score secara manual karena dataset yang digunakan telah memiliki kolom status TB/U dari Puskesmas Pagedangan. Dengan demikian, standar z-score berfungsi untuk menjelaskan dasar konseptual pembentukan kategori status stunting, sedangkan proses pemodelan menggunakan label target yang sudah tersedia pada dataset. Penjelasan ini penting agar penggunaan teori z-score tidak dipahami sebagai proses perhitungan ulang dalam tahap pemodelan [41].

2.2.5 Posyandu

Posyandu merupakan bentuk Upaya Kesehatan Bersumber Daya Masyarakat yang dikelola dari, oleh, untuk, dan bersama masyarakat dengan bimbingan teknis dari Puskesmas [9]. Posyandu berperan sebagai salah satu titik awal pemantauan pertumbuhan balita melalui kegiatan penimbangan, pencatatan, penyuluhan, serta rujukan apabila ditemukan indikasi gangguan pertumbuhan.

Dalam pelaksanaannya, kegiatan Posyandu dilakukan melalui sistem lima langkah, yaitu pendaftaran, penimbangan, pengisian KMS, penyuluhan, dan pelayanan kesehatan [9]. Kader Posyandu berperan dalam membantu pelaksanaan kegiatan tersebut, termasuk melakukan pencatatan hasil penimbangan balita pada Buku KIA atau KMS dan buku register Posyandu.

Apabila ditemukan kondisi tertentu, seperti berat badan tidak naik dua kali berturut-turut atau berada di bawah garis merah, kader wajib melakukan rujukan ke Puskesmas atau Poskesdes untuk mendapatkan tindak lanjut [10].

Peran Posyandu dalam pemantauan pertumbuhan balita menunjukkan bahwa data antropometri yang digunakan dalam layanan kesehatan primer tidak hanya berasal dari proses pemeriksaan di Puskesmas, tetapi juga dapat berasal dari pengukuran dan pencatatan awal di Posyandu. Hal ini menjadikan Posyandu sebagai bagian penting dalam rantai pencatatan data pertumbuhan balita.

2.2.6 Program Makan Bergizi Gratis

Program Makan Bergizi Gratis atau MBG merupakan program intervensi gizi yang ditujukan untuk mendukung pemenuhan gizi kelompok sasaran, termasuk anak balita, ibu hamil, ibu menyusui, dan peserta didik. Dalam konteks kesehatan anak, program ini memiliki keterkaitan dengan upaya pencegahan stunting karena stunting berkaitan dengan kekurangan gizi kronis yang terjadi dalam jangka panjang [8].

MBG tidak dapat dinilai hanya dari pelaksanaan distribusi makanan. Program ini juga memerlukan pemantauan status gizi secara berkala agar perubahan kondisi gizi sasaran dapat diketahui. Pada balita, pemantauan tersebut dapat dilakukan melalui data antropometri, seperti umur, jenis kelamin, berat badan, serta tinggi atau panjang badan. Data tersebut menjadi dasar untuk melihat apakah intervensi gizi berjalan searah dengan perbaikan status pertumbuhan anak.

Dalam penelitian ini, MBG diposisikan sebagai konteks kebijakan yang memperkuat pentingnya pemantauan status stunting berbasis data. Penelitian ini tidak mengevaluasi efektivitas program MBG secara langsung karena data yang digunakan hanya berupa data antropometri balita pada periode Januari 2026. Namun, hasil penelitian dapat menjadi dasar awal untuk menunjukkan bagaimana data antropometri dapat dimanfaatkan dalam klasifikasi awal status

stunting dan mendukung kebutuhan monitoring program gizi pada tingkat layanan kesehatan primer.

2.2.7 Pencatatan Data Antropometri Balita

Data antropometri balita merupakan data hasil pengukuran fisik yang digunakan untuk menilai pertumbuhan dan status gizi anak. Data ini umumnya mencakup usia, jenis kelamin, berat badan, tinggi atau panjang badan, serta status gizi berdasarkan indikator tertentu [44]. Dalam konteks stunting, indikator yang digunakan adalah tinggi badan menurut umur atau panjang badan menurut umur [45].

Pada layanan kesehatan primer, data antropometri dapat diperoleh melalui kegiatan Posyandu dan Puskesmas. Data pertama kali diperoleh dari proses pengukuran di lapangan, kemudian dicatat secara manual pada Buku KIA, KMS, atau register Posyandu [46], [47]. Selanjutnya, apabila terdapat indikasi gangguan pertumbuhan, balita dapat dirujuk ke Puskesmas untuk mendapatkan pemeriksaan lebih lanjut oleh tenaga kesehatan atau ahli gizi [48].

Proses tersebut menunjukkan bahwa data antropometri balita dapat melalui beberapa tahap pencatatan sebelum digunakan dalam sistem pencatatan dan pengolahan data kesehatan. Tahapan yang berlapis tersebut berpotensi menimbulkan ketidaksesuaian data, terutama apabila terdapat kesalahan pengukuran, kesalahan pembacaan ulang, nilai yang tidak lengkap, atau kesalahan input pada sistem pencatatan digital [14] [49].

2.2.8 Kualitas Data Kesehatan

Kualitas data merupakan aspek penting dalam penerapan sistem informasi kesehatan dan machine learning. Data yang tidak berkualitas dapat memengaruhi hasil analisis, menurunkan performa model, serta menyebabkan kesalahan dalam pengambilan keputusan. Dalam konteks data kesehatan, kualitas data umumnya berkaitan dengan kelengkapan, akurasi, konsistensi, validitas, dan ketepatan waktu [16].

Pada data antropometri balita, permasalahan kualitas data dapat muncul dalam berbagai bentuk, seperti nilai kosong, data duplikat, kesalahan satuan,

nilai ekstrem, kesalahan input, atau ketidaksesuaian antara data antropometri dan status gizi yang tercatat. Penelitian terkait aplikasi Posyandu menunjukkan bahwa kualitas data pemantauan balita masih dapat menghadapi tantangan pada aspek akurasi dan konsistensi [12]. Selain itu, penelitian lain menunjukkan bahwa ketidakakuratan pengukuran antropometri oleh kader Posyandu dapat memengaruhi klasifikasi status stunting balita [14].

Dalam penelitian ini, kualitas data menjadi perhatian karena model machine learning tidak hanya digunakan untuk klasifikasi status stunting, tetapi juga diposisikan sebagai alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali. Apabila terdapat perbedaan antara hasil prediksi model dengan status yang tercatat, data tersebut dapat dianggap sebagai data yang berpotensi perlu diperiksa kembali oleh pihak terkait.

2.2.9 Data Cross-Sectional

Data cross-sectional merupakan data yang dikumpulkan pada satu waktu atau satu periode tertentu untuk menggambarkan kondisi objek penelitian pada periode tersebut [50]. Berbeda dengan data time series yang digunakan untuk melihat perubahan dari waktu ke waktu, data cross-sectional digunakan untuk menganalisis hubungan antarvariabel pada satu titik waktu tertentu.

Dalam konteks penelitian ini, dataset yang digunakan merupakan data antropometri balita pada periode Januari 2026. Oleh karena itu, data tersebut diposisikan sebagai data cross-sectional atau data snapshot yang menggambarkan kondisi balita pada periode pengukuran tersebut. Penelitian ini tidak bertujuan untuk memprediksi perubahan status stunting dari waktu ke waktu, melainkan melakukan klasifikasi status stunting berdasarkan kondisi antropometri balita pada periode data yang tersedia.

Penggunaan data satu periode tetap relevan dalam penelitian klasifikasi karena model dibangun untuk mengenali pola hubungan antara variabel input dan label target pada dataset tersebut [51]. Variabel input yang digunakan meliputi umur, jenis kelamin, berat badan, dan tinggi badan atau panjang badan, sedangkan label target diperoleh dari status TB/U. Dengan demikian, data satu

bulan dapat digunakan sebagai dasar pemodelan selama ruang lingkup penelitian dijelaskan sebagai klasifikasi berdasarkan data pada satu periode pengukuran, bukan sebagai analisis tren atau prediksi longitudinal.

2.2.10 *Machine Learning*

Machine learning merupakan cabang dari kecerdasan buatan yang memungkinkan komputer mempelajari pola dari data untuk menghasilkan prediksi atau klasifikasi tanpa diprogram secara eksplisit untuk setiap kasus [52], [53]. Model *machine learning* dibangun berdasarkan data historis, kemudian digunakan untuk mengenali pola pada data baru berdasarkan hubungan antara variabel input dan output.

Secara umum, *machine learning* dapat dibagi menjadi tiga jenis utama, yaitu *supervised learning*, *unsupervised learning*, dan *reinforcement learning*. Penelitian ini menggunakan pendekatan *supervised learning* karena dataset yang digunakan memiliki label status stunting sebagai target klasifikasi. Model dilatih menggunakan data antropometri balita yang telah memiliki label, kemudian digunakan untuk memprediksi kelas status stunting pada data lain [54].

Klasifikasi merupakan salah satu tugas dalam *supervised learning* yang bertujuan untuk mengelompokkan data ke dalam kelas atau kategori tertentu berdasarkan fitur yang tersedia [55]. Dalam penelitian ini, klasifikasi digunakan untuk menentukan status stunting balita berdasarkan data antropometri. Fitur yang digunakan disesuaikan dengan variabel yang tersedia pada dataset Puskesmas, sedangkan label klasifikasi disesuaikan dengan status gizi yang tercatat. Output klasifikasi dapat berupa kelas seperti *Severely Stunted*, *Stunted*, atau *Normal*, sesuai dengan ketersediaan label pada dataset.

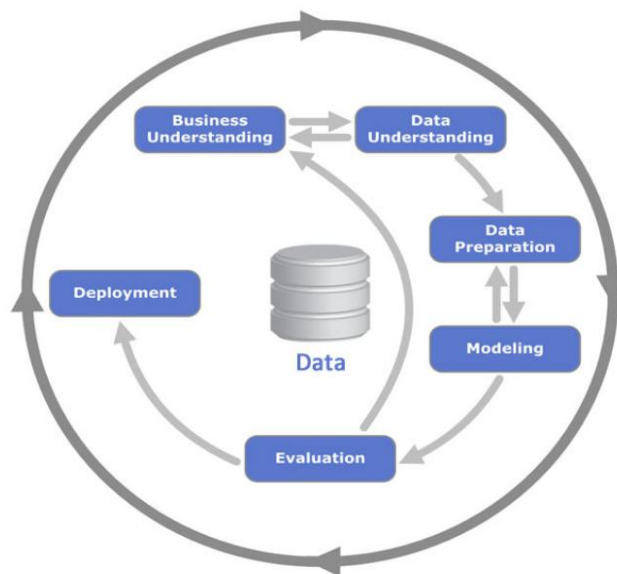
Dalam konteks penelitian ini, *machine learning* tidak hanya digunakan untuk menghasilkan prediksi status stunting, tetapi juga diposisikan sebagai alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali. Apabila hasil prediksi model berbeda dengan status yang tercatat pada dataset, maka data tersebut dapat ditandai sebagai data yang berpotensi perlu

diperiksa kembali. Dengan demikian, klasifikasi dalam penelitian ini tidak hanya berfungsi sebagai metode prediksi, tetapi juga sebagai pendekatan analitik untuk membantu peninjauan awal terhadap data status stunting balita yang berpotensi perlu diperiksa kembali [56].

2.3 Framework/Algoritma yang digunakan

2.3.1 Framework CRISP-DM

Cross-Industry Standard Process for Data Mining atau CRISP-DM merupakan framework yang banyak digunakan dalam penelitian dan proyek data mining. CRISP-DM menyediakan tahapan kerja yang sistematis untuk mengolah data mentah menjadi informasi atau model yang dapat digunakan dalam pengambilan Keputusan [57]. Tahapan dari CRISP-DM dapat dilihat pada Gambar 2.1



Gambar 2.1 Framework CRISP-DM

Sumber: [57]

Dalam penelitian ini, CRISP-DM digunakan untuk mengarahkan seluruh proses penelitian. Tahap *business understanding* digunakan untuk memahami permasalahan penelitian, yaitu kebutuhan klasifikasi status stunting balita dan peninjauan awal terhadap konsistensi data antropometri. Tahap *data*

understanding digunakan untuk memahami struktur, isi, dan karakteristik dataset Puskesmas. Tahap *data preparation* digunakan untuk menyiapkan data agar dapat digunakan dalam proses pemodelan. Tahap *modeling* digunakan untuk membangun model klasifikasi menggunakan algoritma *K-Nearest Neighbor* dan *Support Vector Machine*. Tahap *evaluation* digunakan untuk membandingkan performa kedua algoritma, sedangkan tahap *deployment* diwujudkan dalam bentuk simulasi atau prototipe penggunaan model sebagai alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali [58].

Berikut adalah penjelasan keenam fase CRISP-DM beserta penerapannya dalam penelitian ini:

1) *Business Understanding*

Tahap business understanding merupakan tahap awal dalam CRISP-DM yang bertujuan untuk memahami permasalahan dari sudut pandang domain penelitian. Dalam penelitian ini, permasalahan utama berkaitan dengan pemanfaatan data antropometri balita di Puskesmas untuk mendukung klasifikasi status stunting serta membantu mengevaluasi konsistensi data.

Pada tahap ini, ditentukan tujuan penelitian, yaitu membangun model klasifikasi status stunting balita menggunakan algoritma *KNeighborsClassifier* dan *Support Vector Classifier*, membandingkan performa kedua algoritma, serta membuat simulasi penggunaan model terbaik sebagai alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali. Kriteria keberhasilan penelitian ditentukan berdasarkan hasil evaluasi model menggunakan metrik accuracy, precision, recall, dan F1-score.

2) *Data Understanding*

Tahap data understanding bertujuan untuk memahami karakteristik dataset yang digunakan dalam penelitian. Dataset yang digunakan merupakan data antropometri balita dari Puskesmas Pagedangan. Data

tersebut dapat mencakup variabel seperti usia balita, jenis kelamin, berat badan, tinggi atau panjang badan, serta status gizi atau status stunting.

Pada tahap ini dilakukan eksplorasi awal terhadap dataset, seperti pemeriksaan jumlah data, jumlah variabel, tipe data, distribusi kelas target, nilai kosong, data duplikat, dan kemungkinan nilai ekstrem. Tahap ini penting karena kualitas data akan memengaruhi hasil pemodelan. Apabila ditemukan data yang tidak lengkap, tidak konsisten, atau memiliki nilai tidak wajar, maka data tersebut perlu ditangani pada tahap berikutnya.

3) *Data Preparation*

Tahap data preparation merupakan tahap persiapan data sebelum dilakukan pemodelan. Pada tahap ini, data mentah diolah menjadi dataset yang siap digunakan oleh algoritma machine learning. Aktivitas yang dilakukan dapat meliputi pembersihan data, penanganan nilai kosong, penghapusan data duplikat, penyesuaian format data, encoding variabel kategorikal, serta pemisahan fitur dan label.

Dalam penelitian ini, fitur yang digunakan disesuaikan dengan variabel yang tersedia pada dataset Puskesmas. Variabel seperti usia, jenis kelamin, berat badan, dan tinggi atau panjang badan digunakan sebagai fitur input, sedangkan status stunting atau status gizi digunakan sebagai label target. Apabila terdapat variabel kategorikal seperti jenis kelamin, maka variabel tersebut perlu dikonversi ke bentuk numerik agar dapat diproses oleh algoritma machine learning.

Tahap ini juga dapat mencakup pembagian data menjadi data latih dan data uji. Data latih digunakan untuk membangun model, sedangkan data uji digunakan untuk mengukur performa model pada data yang belum pernah dilihat sebelumnya. Selain itu, penelitian ini juga menggunakan k-fold cross-validation untuk memperoleh evaluasi yang lebih stabil.

4) *Modeling*

Tahap modeling merupakan tahap pembangunan model machine learning. Pada tahap ini, algoritma KNeighborsClassifier dan Support Vector Classifier diterapkan untuk melakukan klasifikasi status stunting balita. Kedua algoritma dilatih menggunakan data latih yang telah melalui tahap persiapan data.

Model K-Nearest Neighbor dibangun dengan menentukan jumlah tetangga terdekat atau nilai K yang digunakan dalam proses klasifikasi. Sementara itu, model Support Vector Machine dibangun dengan menentukan kernel dan parameter yang digunakan untuk membentuk hyperplane pemisah antar kelas. Hasil dari tahap ini adalah model klasifikasi yang mampu memprediksi status stunting balita berdasarkan data antropometri.

5) *Evaluation*

Tahap *evaluation* digunakan untuk menilai performa model yang telah dibangun. Evaluasi dilakukan dengan membandingkan hasil prediksi model terhadap label aktual pada dataset. Dalam penelitian ini, evaluasi dilakukan menggunakan confusion matrix dan metrik evaluasi seperti accuracy, precision, recall, dan F1-score.

Selain itu, hasil evaluasi *K-Nearest Neighbor* dan *Support Vector Machine* dibandingkan untuk menentukan algoritma dengan performa terbaik. Model terbaik tidak hanya dilihat dari nilai akurasi, tetapi juga dari kemampuan model dalam mengenali kelas stunting secara tepat. Oleh karena itu, recall dan F1-score menjadi metrik penting karena berkaitan dengan kemampuan model mendeteksi kasus stunting dan menjaga keseimbangan antara ketepatan dan kelengkapan prediksi.

6) *Deployment*

Tahap *deployment* merupakan tahap penerapan hasil model ke dalam bentuk yang dapat digunakan atau disimulasikan. Dalam penelitian ini, deployment tidak dilakukan dalam bentuk integrasi langsung dengan

sistem Puskesmas atau SIGIZI KESGA. Deployment diwujudkan dalam bentuk website yang dapat menampilkan hasil klasifikasi status stunting berdasarkan data input pengguna.

Simulasi tersebut bertujuan untuk menggambarkan bagaimana model terbaik dapat digunakan sebagai alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali. Pengguna dapat memasukkan data antropometri balita, kemudian sistem menampilkan hasil klasifikasi status stunting. Apabila hasil prediksi model berbeda dengan status yang tercatat pada dataset, data tersebut dapat ditandai sebagai data yang perlu diperiksa kembali oleh tenaga kesehatan atau pihak terkait [58].

2.3.2 *K-Nearest Neighbor*

K-Nearest Neighbor atau KNN merupakan algoritma supervised learning yang dapat digunakan untuk tugas klasifikasi. Algoritma ini bekerja dengan menentukan kelas suatu data baru berdasarkan kedekatannya dengan sejumlah data latih terdekat [59]. Jumlah tetangga terdekat yang digunakan ditentukan oleh nilai K.

Pada proses klasifikasi, KNN menghitung jarak antara data baru dengan seluruh data latih. Setelah jarak dihitung, algoritma memilih sebanyak K data terdekat. Kelas yang paling banyak muncul di antara tetangga terdekat tersebut akan menjadi hasil prediksi untuk data baru [60].

Dalam konteks penelitian ini, KNN digunakan untuk mengklasifikasikan status stunting balita berdasarkan umur, jenis kelamin, berat badan, dan tinggi badan atau panjang badan. Misalnya, apabila data balita baru memiliki karakteristik yang paling dekat dengan data-data balita yang berstatus stunted, maka model KNN dapat mengklasifikasikan balita tersebut ke dalam kelas stunted.

Kelebihan KNN adalah konsepnya sederhana dan mudah dipahami. Namun, KNN memiliki kelemahan karena sangat bergantung pada skala fitur dan

pemilihan nilai K. Oleh karena itu, dalam penelitian ini dilakukan standardization atau scaling agar fitur numerik umur, berat badan, dan tinggi atau panjang badan berada pada skala yang sebanding [61].

2.3.3 Euclidean Distance

Euclidean Distance merupakan salah satu metode pengukuran jarak yang umum digunakan dalam algoritma KNN. Jarak ini digunakan untuk menghitung kedekatan antara satu data dengan data lainnya berdasarkan nilai fitur yang dimiliki.

Rumus Euclidean Distance adalah sebagai berikut:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (2.1)$$

Keterangan :

- 1) $d(x, y)$ adalah jarak antara x dan data y
- 2) x_i adalah nilai fitur ke- i pada data pertama.
- 3) y_i adalah nilai fitur ke- i pada data kedua.
- 4) n adalah jumlah fitur yang digunakan.

Dalam penelitian ini, Euclidean Distance dapat digunakan untuk menghitung jarak antar data balita berdasarkan fitur umur, berat badan, dan tinggi atau panjang badan setelah dilakukan standardization. Penggunaan standardization penting karena apabila fitur memiliki skala yang berbeda, fitur dengan nilai numerik lebih besar dapat mendominasi hasil perhitungan jarak.

2.3.4 Pemilihan Nilai K pada KNN

Nilai K merupakan parameter penting dalam algoritma KNN karena menentukan jumlah tetangga terdekat yang digunakan dalam proses klasifikasi. Nilai K yang terlalu kecil dapat membuat model terlalu sensitif terhadap noise pada data. Sebaliknya, nilai K yang terlalu besar dapat membuat model terlalu umum dan mengurangi kemampuan model dalam menangkap pola lokal pada data.

Pemilihan nilai K dapat dilakukan melalui pengujian beberapa nilai K, misalnya $K = 3, 5, 7, 9$, dan seterusnya. Setiap nilai K dievaluasi menggunakan metrik accuracy, precision, recall, dan F1-score. Dalam penelitian ini, pemilihan nilai K dilakukan dengan mempertimbangkan hasil evaluasi melalui k-fold cross-validation. Nilai K terbaik dipilih berdasarkan performa yang paling optimal, terutama pada metrik recall dan F1-score.

Recall menjadi penting karena dalam konteks klasifikasi stunting, kesalahan berupa false negative dapat berdampak serius. False negative terjadi ketika balita yang sebenarnya mengalami stunting diprediksi sebagai normal, sehingga berisiko tidak mendapatkan perhatian atau tindak lanjut yang sesuai.

2.3.5 Support Vector Machine

Support Vector Machine atau SVM merupakan algoritma supervised learning yang dapat digunakan untuk klasifikasi. SVM bekerja dengan mencari batas pemisah terbaik atau hyperplane yang mampu memisahkan data antar kelas dengan margin terbesar [62]. Margin adalah jarak antara hyperplane dan data terdekat dari masing-masing kelas. Data terdekat yang memengaruhi pembentukan hyperplane disebut support vector.

Dalam klasifikasi biner, SVM mencari hyperplane yang memisahkan dua kelas. Namun, dalam penelitian ini label target terdiri dari tiga kelas, yaitu severely stunted, stunted, dan normal, sehingga SVM digunakan dalam pendekatan multi-kelas. Pendekatan multi-kelas pada SVM umumnya dapat dilakukan dengan strategi one-vs-rest atau one-vs-one.

Dalam konteks penelitian ini, SVM digunakan untuk mengklasifikasikan status stunting balita berdasarkan data antropometri. SVM relevan digunakan karena mampu menangani data dengan pola yang kompleks melalui penggunaan kernel. Namun, seperti KNN, SVM juga sensitif terhadap skala data, sehingga tahap standardization atau scaling menjadi bagian penting dalam preprocessing [63].

2.3.6 Hyperplane dan Margin

Hyperplane merupakan batas keputusan yang digunakan oleh SVM untuk memisahkan data ke dalam kelas yang berbeda. Pada ruang dua dimensi, hyperplane dapat digambarkan sebagai garis. Pada ruang tiga dimensi, hyperplane dapat berupa bidang. Pada dimensi yang lebih tinggi, hyperplane menjadi batas pemisah dalam ruang fitur berdimensi tinggi.

Margin adalah jarak antara hyperplane dengan data terdekat dari masing-masing kelas [62]. Tujuan utama SVM adalah mencari hyperplane dengan margin terbesar agar pemisahan antar kelas menjadi lebih optimal. Semakin besar margin yang dihasilkan, semakin baik kemampuan model dalam melakukan generalisasi pada data baru.

Dalam konteks klasifikasi status stunting, hyperplane digunakan untuk memisahkan pola data antropometri antar kelas, seperti normal, stunted, dan severely stunted. Karena data kesehatan sering kali memiliki pola yang tidak sepenuhnya linear, penggunaan kernel pada SVM dapat membantu memetakan data ke ruang fitur yang lebih sesuai [63].

2.3.7 Kernel pada SVM

Kernel merupakan fungsi yang digunakan dalam SVM untuk memetakan data dari ruang fitur awal ke ruang fitur berdimensi lebih tinggi. Tujuan penggunaan kernel adalah agar data yang tidak dapat dipisahkan secara linear pada ruang awal dapat menjadi lebih mudah dipisahkan pada ruang fitur baru [63].

Beberapa jenis kernel yang umum digunakan dalam SVM adalah:

1. Linear Kernel, digunakan ketika data dapat dipisahkan secara linear.
2. Polynomial Kernel, digunakan untuk menangkap pola hubungan non-linear berbentuk polinomial.
3. Radial Basis Function Kernel, digunakan untuk menangani pola non-linear yang lebih kompleks.

4. Sigmoid Kernel, digunakan pada kondisi tertentu yang menyerupai fungsi aktivasi pada jaringan saraf [64].

Dalam penelitian ini, penggunaan kernel dapat disesuaikan dengan hasil pengujian model. Kernel yang umum digunakan untuk klasifikasi data dengan pola non-linear adalah Radial Basis Function [65]. Namun, pemilihan kernel tetap perlu diuji melalui evaluasi model agar sesuai dengan karakteristik dataset Puskesmas Pagedangan.

2.3.8 *Standardization* atau *Scaling*

Standardization atau scaling merupakan tahap preprocessing yang bertujuan untuk menyamakan skala antar fitur numerik. Tahap ini penting karena algoritma KNN dan SVM sensitif terhadap perbedaan skala data. umur, berat badan, dan tinggi atau panjang badan memiliki satuan dan rentang nilai yang berbeda. Apabila tidak dilakukan scaling, fitur dengan rentang nilai lebih besar dapat memberikan pengaruh yang lebih dominan terhadap model [66], [67].

Salah satu metode standardization yang umum digunakan adalah StandardScaler, yaitu mengubah data sehingga memiliki rata-rata 0 dan standar deviasi 1. Rumus standardization adalah sebagai berikut:

$$z = \frac{x - \mu}{\sigma} \quad (2.2)$$

Keterangan :

- 1) z adalah nilai hasil standardization.
- 2) x adalah nilai asli.
- 3) μ adalah rata-rata fitur.
- 4) σ adalah standar deviasi fitur [68].

Dalam penelitian ini, standardization diterapkan pada fitur numerik seperti umur, berat badan, dan tinggi atau panjang badan. Proses ini dilakukan setelah data dibersihkan dan sebelum model KNN dan SVM dilatih. Penerapan standardization penting agar perhitungan jarak pada KNN dan pembentukan

hyperplane pada SVM tidak dipengaruhi secara berlebihan oleh perbedaan skala antar fitur [61], [69].

2.3.9 K-Fold Cross-Validation

K-fold cross-validation merupakan teknik validasi model yang digunakan untuk mengevaluasi performa model secara lebih stabil dan objektif [70]. Teknik ini dilakukan dengan membagi dataset menjadi sejumlah bagian yang disebut *fold*. Model kemudian dilatih dan diuji secara berulang sebanyak jumlah *fold* yang digunakan. Pada setiap iterasi, satu *fold* digunakan sebagai data uji, sedangkan *fold* lainnya digunakan sebagai data latih.

Secara umum, jika dataset dibagi menjadi K bagian, maka proses pelatihan dan pengujian dilakukan sebanyak K kali. Setiap bagian data akan menjadi data uji tepat satu kali, sedangkan bagian lainnya digunakan sebagai data latih. Nilai performa akhir diperoleh dengan menghitung rata-rata hasil evaluasi dari seluruh iterasi. Rumus rata-rata performa *k-fold cross-validation* dapat dituliskan sebagai berikut:

$$CV = \frac{1}{K} \sum_{i=1}^K M_i \quad (2.3)$$

Keterangan :

- 1) CV adalah nilai rata-rata performa hasil *cross-validation*
- 2) K adalah jumlah *fold*.
- 3) Nilai metrik evaluasi pada iterasi ke- i : M_i
- 4) Indeks iterasi pengujian : i

Sebagai contoh, apabila digunakan $K=5$, maka dataset dibagi menjadi lima bagian. Pada iterasi pertama, bagian pertama digunakan sebagai data uji dan empat bagian lainnya digunakan sebagai data latih. Pada iterasi kedua, bagian kedua digunakan sebagai data uji dan bagian lainnya sebagai data latih. Proses ini dilakukan hingga seluruh bagian pernah menjadi data uji. Jika nilai F1-score dari lima iterasi adalah 0,82, 0,84, 0,81, 0,85, dan 0,83, maka nilai rata-rata F1-score adalah jumlah seluruh nilai dibagi lima, yaitu 0,83.

Kelebihan k-fold cross-validation adalah mampu mengurangi ketergantungan hasil evaluasi terhadap satu kali pembagian data latih dan data uji [71], [72]. Teknik ini juga membuat seluruh data memiliki kesempatan untuk digunakan sebagai data latih maupun data uji. Dengan demikian, hasil evaluasi menjadi lebih representatif dan stabil. Hal ini penting dalam penelitian klasifikasi stunting karena distribusi data antar kelas dapat memengaruhi hasil model.

Namun, k-fold cross-validation juga memiliki kekurangan. Teknik ini membutuhkan waktu komputasi yang lebih besar karena proses pelatihan dan pengujian dilakukan berulang kali [70]. Semakin besar nilai K, semakin banyak proses pelatihan yang perlu dilakukan. Selain itu, apabila dataset memiliki distribusi kelas yang tidak seimbang, pembagian fold secara acak dapat menyebabkan beberapa fold memiliki proporsi kelas yang tidak representatif.

Untuk mengatasi permasalahan distribusi kelas yang tidak seimbang, dapat digunakan stratified k-fold cross-validation [73]. Teknik ini menjaga proporsi kelas pada setiap fold agar tetap menyerupai distribusi kelas pada dataset asli. Dalam konteks klasifikasi status stunting, pendekatan ini penting karena jumlah data pada kelas tertentu, seperti Severely Stunted, dapat lebih sedikit dibandingkan kelas Normal. Dengan menjaga proporsi kelas pada setiap fold, evaluasi model dapat menjadi lebih adil dan representatif.

Dalam penelitian ini, k-fold cross-validation digunakan untuk mengevaluasi performa KNeighborsClassifier dan Support Vector Classifier. Teknik ini membantu memastikan bahwa hasil perbandingan kedua algoritma tidak hanya bergantung pada satu skenario pembagian data, melainkan diperoleh dari beberapa kali proses pelatihan dan pengujian.

2.3.10 Confusion Matrix dan Metrik Evaluasi

Confusion matrix merupakan alat evaluasi yang digunakan untuk menilai performa model klasifikasi dengan membandingkan hasil prediksi model terhadap label aktual. Melalui confusion matrix, dapat diketahui jumlah prediksi yang benar maupun salah pada setiap kelas. Pada klasifikasi biner, terdapat

empat komponen utama dalam confusion matrix, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN) [74], [75].

True Positive menunjukkan jumlah data positif yang diprediksi positif dengan benar. True Negative menunjukkan jumlah data negatif yang diprediksi negatif dengan benar. False Positive menunjukkan jumlah data negatif yang salah diprediksi sebagai positif. Sementara itu, False Negative menunjukkan jumlah data positif yang salah diprediksi sebagai negatif. Dalam konteks klasifikasi stunting, False Negative menjadi perhatian penting karena menunjukkan kondisi ketika balita yang sebenarnya mengalami stunting justru diprediksi sebagai bukan stunting.

Dalam penelitian ini, klasifikasi dapat dilakukan pada lebih dari dua kelas, misalnya Severely Stunted, Stunted, dan Normal, sesuai dengan label yang tersedia pada dataset. Pada klasifikasi multi-kelas, confusion matrix tetap digunakan untuk melihat kesalahan prediksi antar kelas. Metrik evaluasi dapat dihitung untuk setiap kelas, kemudian dirata-ratakan menggunakan pendekatan tertentu, seperti macro average atau weighted average [76], [77]. Pendekatan macro average menghitung rata-rata metrik setiap kelas dengan bobot yang sama, sedangkan weighted average memberikan bobot berdasarkan jumlah data pada masing-masing kelas.

Metrik evaluasi yang digunakan dalam penelitian ini meliputi accuracy, precision, recall, dan F1-score. Accuracy mengukur proporsi seluruh prediksi yang benar terhadap total data. Rumus accuracy adalah sebagai berikut:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2.4)$$

Keterangan :

- 1) Jumlah data positif yang diprediksi positif dengan benar : TP
- 2) Jumlah data negatif yang diprediksi negatif dengan benar : TN
- 3) Jumlah data negatif yang salah diprediksi positif : FP
- 4) Jumlah data positif yang salah diprediksi negatif : FN

Precision mengukur ketepatan model ketika memprediksi suatu kelas sebagai positif. Rumus *precision* adalah sebagai berikut:

$$Precision = \frac{TP}{TP+FP} \quad (2.5)$$

Keterangan :

- 1) Jumlah prediksi positif yang benar : *TP*
- 2) Jumlah prediksi positif yang salah : *FP*

Recall mengukur kemampuan model dalam menemukan seluruh data yang benar-benar termasuk kelas positif. Rumus *recall* adalah sebagai berikut:

$$Recall = \frac{TP}{TP+FN} \quad (2.6)$$

Keterangan :

- 1) Jumlah data positif yang berhasil diprediksi dengan benar : *TP*
- 2) Jumlah data positif yang salah diprediksi sebagai negatif : *FN*

F1-score merupakan metrik yang menggabungkan *precision* dan *recall* ke dalam satu nilai. Rumus *F1-score* adalah sebagai berikut:

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.7)$$

Keterangan :

- 1) Nilai ketepatan prediksi positif : *Precision*
- 2) Kemampuan model dalam menemukan data positif yang sebenarnya : *Recall*

Istilah akurasi digunakan dalam dua konteks yang berbeda. Akurasi data berkaitan dengan ketepatan dan konsistensi data antropometri yang dicatat, seperti umur, berat badan, tinggi badan, serta status TB/U. Sementara itu, accuracy model merupakan metrik evaluasi machine learning yang menunjukkan proporsi prediksi benar terhadap seluruh data uji. Dengan

demikian, accuracy model tidak dapat disamakan langsung dengan akurasi pengukuran antropometri di lapangan.

Dalam konteks klasifikasi stunting, accuracy berguna untuk melihat performa model secara umum. Namun, accuracy saja tidak cukup apabila distribusi kelas tidak seimbang [78], [79]. Sebagai contoh, apabila jumlah balita Normal jauh lebih banyak dibandingkan balita Stunted, model dapat memperoleh nilai accuracy tinggi hanya dengan sering memprediksi kelas Normal. Oleh karena itu, metrik precision, recall, dan F1-score diperlukan untuk memberikan gambaran performa yang lebih lengkap.

Recall menjadi metrik yang penting dalam konteks kesehatan karena berkaitan dengan kemampuan model dalam mendeteksi balita yang benar-benar mengalami stunting. Kesalahan berupa False Negative, yaitu balita stunting yang diprediksi sebagai Normal, dapat berdampak lebih serius karena balita tersebut berpotensi tidak mendapatkan perhatian atau tindak lanjut [80] [81]. Sementara itu, precision penting untuk memastikan bahwa prediksi stunting yang diberikan model memiliki ketepatan yang baik. F1-score digunakan untuk menilai keseimbangan antara precision dan recall.

Dalam penelitian ini, hasil confusion matrix dan metrik evaluasi digunakan untuk membandingkan performa KNeighborsClassifier dan Support Vector Classifier. Model dengan performa terbaik kemudian digunakan sebagai dasar dalam simulasi atau prototipe alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali. Apabila hasil prediksi model berbeda dengan status stunting yang tercatat, data tersebut dapat ditandai sebagai data yang berpotensi perlu diperiksa kembali oleh tenaga kesehatan atau pihak terkait.

2.4 Tools/Software

2.4.1 Python



Gambar 2.2 Logo Python

Python adalah bahasa pemrograman tingkat tinggi yang bersifat *interpreted* dan berorientasi objek, serta telah menjadi standar *de facto* dalam bidang *data science* dan *machine learning*. Python dipilih dalam penelitian ini karena memiliki ekosistem *library* yang sangat lengkap untuk seluruh tahapan *data mining*, mulai dari pengolahan data, pemodelan, hingga visualisasi hasil dalam satu lingkungan kerja yang terintegrasi. Selain memiliki tingkat keterbacaan kode yang tinggi, Python juga didukung oleh komunitas yang besar sehingga dokumentasi dan referensi pemecahan masalah mudah ditemukan. Versi yang digunakan dalam penelitian ini adalah Python 3.x [82].

2.4.2 Jupyter Notebook



Gambar 2.3 Logo Jupyter

Jupyter Notebook adalah aplikasi web *open-source* yang memungkinkan pembuatan dan berbagi dokumen yang memuat kode yang dapat dieksekusi,

persamaan, visualisasi, dan teks naratif dalam satu file berekstensi *.ipynb*. Jupyter Notebook dipilih sebagai lingkungan pengembangan utama dalam penelitian ini karena mendukung eksekusi kode secara interaktif sel per sel, memudahkan eksplorasi data secara bertahap, serta mengintegrasikan kode Python, output hasil, dan dokumentasi analisis dalam satu dokumen yang runtut dan mudah dipresentasikan. Jupyter Notebook dijalankan secara lokal menggunakan distribusi Anaconda, yang menyediakan Python beserta seluruh *library data science* yang dibutuhkan dalam satu paket instalasi [82].

2.4.3 *scikit-learn*

Scikit-learn merupakan *library* machine learning berbasis Python yang menyediakan implementasi algoritma K-Nearest Neighbor melalui *KNeighborsClassifier* dan Support Vector Machine melalui *SVC*. Selain itu, scikit-learn juga menyediakan utilitas preprocessing seperti *StandardScaler*, validasi model seperti *StratifiedKFold* dan *cross_val_score*, serta evaluasi model seperti *classification_report* dan *confusion_matrix* [4], [30], [83].

2.4.4 *pandas* dan *NumPy*

Pandas merupakan *library* Python yang digunakan untuk manipulasi dan analisis data tabular, serta menjadi fondasi dalam pengolahan dataset pada penelitian ini. Pandas digunakan untuk membaca dataset berformat *.csv*, melakukan eksplorasi data awal, menangani nilai hilang, melakukan *encoding* variabel kategorikal, serta transformasi data. Sementara itu, NumPy digunakan untuk operasi array dan komputasi numerik yang menjadi dasar bagi sebagian besar *library data science* Python, termasuk operasi matriks dan perhitungan statistik [82].

2.4.5 *matplotlib* dan *seaborn*

Matplotlib dan Seaborn merupakan *library* visualisasi data yang digunakan secara komplementer dalam penelitian ini. Matplotlib digunakan untuk membuat grafik dasar seperti histogram distribusi variabel, sedangkan Seaborn digunakan untuk visualisasi yang lebih informatif seperti *heatmap* korelasi antar variabel dan perbandingan performa model. Visualisasi yang dihasilkan

mendukung analisis pada fase *Data Understanding* serta interpretasi hasil pada fase *Evaluation* dalam framework CRISP-DM [82].

2.4.6 Flask

Flask merupakan framework web berbasis Python yang digunakan untuk membangun aplikasi web secara ringan dan fleksibel. Flask termasuk microframework karena menyediakan komponen inti untuk pengembangan web, seperti routing, request handling, response handling, dan template rendering, tanpa memaksakan struktur proyek tertentu [84]. Dengan karakteristik tersebut, Flask dapat digunakan untuk membangun prototipe aplikasi web yang terhubung dengan model machine learning [85].

Dalam pengembangan aplikasi machine learning, Flask dapat berperan sebagai backend yang menerima input dari pengguna, memproses data input, menjalankan model yang telah disimpan, dan mengembalikan hasil prediksi ke antarmuka pengguna. Flask juga dapat diintegrasikan dengan library Python lain seperti pandas, joblib, dan scikit-learn. Integrasi tersebut memungkinkan model machine learning yang telah dilatih sebelumnya digunakan kembali dalam aplikasi web [86].

Pada penelitian ini, Flask digunakan untuk membangun prototipe website klasifikasi awal status stunting balita. Aplikasi menerima input berupa umur, jenis kelamin, berat badan, dan tinggi atau panjang badan. Data input kemudian diproses menggunakan scaler dan model SVM yang telah disimpan dalam format file model. Hasil prediksi model dikembalikan dalam bentuk output klasifikasi berupa normal, stunted, atau severely stunted. Prototipe ini tidak digunakan sebagai sistem diagnosis medis, melainkan sebagai alat bantu klasifikasi awal dan alat bantu informasi awal untuk menandai data yang berpotensi perlu diperiksa kembali [87].

2.4.7 Visual Studio Code

Visual Studio Code merupakan code editor yang digunakan untuk menulis, mengelola, dan menjalankan kode program. Visual Studio Code mendukung berbagai bahasa pemrograman, termasuk Python, serta menyediakan fitur

seperti syntax highlighting, debugging, terminal terintegrasi, manajemen folder proyek, dan dukungan extension [88]. Fitur-fitur tersebut membantu proses pengembangan program menjadi lebih terstruktur dan efisien [89].

Dalam pengembangan aplikasi machine learning, Visual Studio Code dapat digunakan untuk menulis kode preprocessing, pemodelan, evaluasi model, serta pengembangan prototipe aplikasi web. Selain itu, terminal terintegrasi pada Visual Studio Code memudahkan pengguna dalam menjalankan file Python, menginstal library, serta menguji aplikasi secara lokal [90].

Pada penelitian ini, Visual Studio Code digunakan sebagai lingkungan pengembangan untuk membuat dan mengelola kode program. Kode yang dikembangkan mencakup proses pemanggilan model, pembuatan antarmuka aplikasi, pengolahan input pengguna, serta penampilan hasil klasifikasi. Dengan menggunakan Visual Studio Code, proses pengembangan prototipe dapat dilakukan secara lebih rapi karena seluruh file proyek dapat dikelola dalam satu workspace [91].

