

**PENERAPAN METODE *SUPPORT VECTOR MACHINE* DENGAN PENDEKATAN
MULTIPLE KERNEL LEARNING UNTUK KLASIFIKASI STUNTING
BERBASIS WEB**

SKRIPSI

Oleh:

MUHAMMAD SALMAN ALFARIZY YONDA PUTRA
NIM. 220605110061



**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

**PENERAPAN METODE SUPPORT VECTOR MACHINE DENGAN
PENDEKATAN MULTIPLE KERNEL LEARNING UNTUK
KLASIFIKASI STUNTING BERBASIS WEB**

SKRIPSI

Diajukan Kepada:

Universitas Islam Negeri Maulana Malik Ibrahim Malang
Untuk Memenuhi Salah Satu Persyaratan dalam
Memperoleh Gelar Sarjana Komputer (S.Kom)

Oleh :

MUHAMMAD SALMAN ALFARIZY YONDA PUTRA
NIM. 220605110061

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

HALAMAN PERSETUJUAN

PENERAPAN METODE *SUPPORT VECTOR MACHINE* DENGAN PENDEKATAN *MULTIPLE KERNEL LEARNING* UNTUK KLASIFIKASI STUNTING BERBASIS WEB

SKRIPSI

Oleh:

MUHAMMAD SALMAN ALFARIZY YONDA PUTRA
NIM. 220605110061

Telah Diperiksa dan Disetujui untuk Diuji:
Tanggal: 12 Juni 2026

Pembimbing I,



Nurizal Dwi Priandani, M. Kom
NIP. 19920830 202203 1 001

Pembimbing II,



Hani Nurhayati, M.T
NIP. 19780625 200801 2 006

Mengetahui,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyono, M.Kom
NIP. 19841010 201903 1 012

HALAMAN PENGESAHAN

PENERAPAN METODE *SUPPORT VECTOR MACHINE* DENGAN PENDEKATAN *MULTIPLE KERNEL LEARNING* UNTUK KLASIFIKASI STUNTING BERBASIS WEB

SKRIPSI

Oleh:

MUHAMMAD SALMAN ALFARIZY YONDA PUTRA
NIM. 220605110061

Telah Dipertahankan di Depan Dewan Penguji Skripsi
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan
Untuk Memperoleh Gelar Sarjana Komputer (S.Kom)
Tanggal: 23 Juni 2026

Susunan Dewan Penguji:

Ketua Penguji : Dr. Zainal Abidin, M.Kom
NIP. 19760613 200501 1 004



Anggota Penguji I : Fajar Rohman Hariri, M.Kom
NIP. 19890515 201801 1 001

()

Anggota Penguji II : Nurizal Dwi Priandani, M. Kom
NIP. 19920830 202203 1 001

()

Anggota Penguji III : Hani Nurhayati, M.T
NIP. 19780625 200801 2 006

()

Mengetahui dan Mengesahkan,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyono, M.Kom

NIP. 19841010 201903 1 012

PERNYATAAN KEASLIAN TULISAN

Saya yang bertanda tangan di bawah ini:

Nama : Muhammad Salman Alfarizy Yonda Putra
NIM : 220605110061
Fakultas / Program Studi : Sains dan Teknologi / Teknik Informatika
Judul Skripsi : Penerapan Metode Support Vector Machine
dengan Pendekatan Multiple Kernel Learning
untuk Klasifikasi Stunting Berbasis Web

Menyatakan dengan sebenarnya bahwa Skripsi yang saya tulis ini benar-benar merupakan hasil karya saya sendiri, bukan merupakan pengambil alihan data, tulisan, atau pikiran orang lain yang saya akui sebagai hasil tulisan atau pikiran saya sendiri, kecuali dengan mencantumkan sumber cuplikan pada daftar pustaka.

Apabila dikemudian hari terbukti atau dapat dibuktikan skripsi ini merupakan hasil jiplakan, maka saya bersedia menerima sanksi atas perbuatan tersebut.

Malang, 23 Juni 2026

Yang membuat pernyataan,



Muhammad Salman Alfarizy Y.P.
NIM. 220605110061

MOTTO

“Selesaikan Apa Yang Sudah Kamu Mulai”

HALAMAN PERSEMBAHAN

Alhamdulillah *rabbil 'alamin* puji syukur kepada Allah SWT karena berkat rahmat dan petunjuk-Nya penulis dapat menyelesaikan skripsi ini. Shalawat serta salam selalu tercurahkan kepada Rasulullah SAW yang telah membawa kita dari zaman jahiliyah menuju addinul Islam.

Penulis mempersembahkan karya ini kepada diri orang tua, diri sendiri, kakak, para dosen, teman-teman, serta orang-orang yang telah membantu, mendoakan, serta menyemangati penulis dalam menuntaskan skripsi ini.

KATA PENGANTAR

Assalamu 'alaikum warahmatullahi wabarakatuh

Segala puji dan rasa syukur penulis panjatkan ke hadirat Allah *Subhanahu wa Ta'ala* atas limpahan rahmat dan hidayah-Nya, sehingga penulis dapat menyelesaikan skripsi ini dengan baik. Shalawat dan salam semoga senantiasa tercurah kepada Nabi Muhammad SAW, semoga kita semua kelak memperoleh syafaat beliau di hari kiamat. Aamiin.

Penulis ucapkan terima kasih yang sebesar-besarnya kepada seluruh pihak yang memberikan dukungan dan motivasi kepada penulis sehingga dapat menyelesaikan skripsi ini. Ucapan terima kasih ini penulis disampaikan kepada:

1. Prof. Dr. Hj. Ilfi Nur Diana, M.Si, selaku Rektor Universitas Islam Negeri Maulana Malik Ibrahim Malang.
2. Dr. H. Agus Mulyono, M.Si, selaku Dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.
3. Supriyono, M.Kom, selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang.
4. Nurizal Dwi Priandani, M.Kom, selaku dosen pembimbing 1 serta mentor yang telah sangat sabar memberikan bimbingan, saran, dan dukungan yang sangat berarti bagi penulis selama proses penelitian.
5. Hani Nurhayati, M.Kom, selaku dosen pembimbing 2, atas arahan dan pencerahan yang sangat membantu dalam menyempurnakan karya ini.

6. Dr. Zainal Abidin, M.Kom, selaku dosen Penguji I dan Fajar Rohman Hariri, M.Kom, selaku dosen Penguji II yang telah memberikan kritik dan saran yang membangun dalam penyusunan skripsi ini.
7. Seluruh staf dan dosen Program Studi Teknik Informatika yang telah memberikan ilmu, dukungan, dan fasilitas kepada penulis selama masa studi.
8. Orang tua penulis, Almarhum bapak tersayang (Bapak Trisno Budi Wiyono), Ibu tersayang (Ibu Endang Ikhtiariaty), Ayah (Bapak Bambang Wahyu Purnomo), Bunda (Ibu Sri Purwarini), serta Kakak Ismailia Wienda Yasmin Pratama Putri, yang menjadi dasar dan kekuatan terbesar dalam penyelesaian pendidikan sarjana penulis, serta telah memberikan segalanya yang terbaik kepada penulis mulai dari doa, kepercayaan, dan support-nya.
9. Teman-teman seperjuangan, terutama Lefa, Hilmi, Qolbi, Yasril, Singgi, Ghiffari, Ayu, Tiara, Fairuz, Agam, Zaka, dan Epen, serta teman teman lainnya yang senantiasa menjadi rekan berdiskusi dan bertukar pemikiran selama masa perkuliahan hingga proses penyelesaian skripsi ini. Terima kasih atas dukungan, semangat, motivasi, serta berbagai pandangan dan pengalaman yang telah dibagikan, baik dalam bidang akademik maupun kehidupan sehari-hari, sehingga memberikan banyak pelajaran dan dorongan bagi penulis dalam menyelesaikan studi ini.
10. Teman-teman alumni SMK Telkom Malang angkatan 28, terutama Faiz dan Steven. Terima kasih atas kebersamaan, dukungan, semangat, serta berbagai proses suka dan duka yang telah dilalui bersama selama masa

perkuliahan. Semoga segala kebaikan dan pengalaman yang telah dilalui menjadi langkah menuju kesuksesan bagi kita semua.

11. Kepada kekasih tercinta yang telah memberikan motivasi dan semangat, dan bawel dalam menyuruh untuk membuat skripsi ini sehingga skripsi ini bisa dibuat dan diselesaikan dengan semangat yang tinggi dan selesai pada waktunya.

12. Seluruh pihak yang telah terlibat, baik secara langsung maupun tidak langsung dari awal perkuliahan hingga akhir penulisan skripsi ini

Penulis berharap semoga skripsi ini dapat memberikan manfaat di masa yang akan datang. Penulis juga menyadari bahwa karya ini masih memiliki kekurangan. Oleh karena itu, penulis sangat terbuka terhadap segala bentuk masukan dan saran yang membangun untuk pengembangan lebih lanjut.

Wassalamu'alaikum warahmatullahi wabarakatuh.

Malang, 23 Juni 2026

Penulis

DAFTAR ISI

HALAMAN PERSETUJUAN	iii
HALAMAN PENGESAHAN.....	iv
PERNYATAAN KEASLIAN TULISAN	v
MOTTO	vi
HALAMAN PERSEMBAHAN	vii
KATA PENGANTAR.....	viii
DAFTAR ISI.....	xi
DAFTAR GAMBAR.....	xiii
DAFTAR TABEL	xiv
ABSTRAK	xv
ABSTRACT.....	xvi
مستخلص البحث	xvii
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	5
1.3 Tujuan Penelitian	6
1.4 Manfaat Penelitian	6
1.5 Batasan Penelitian.....	7
BAB II STUDI PUSTAKA	8
2.1 Penelitian Terkait	8
2.2 Stunting.....	11
2.3 <i>Support Vector Machine</i> (SVM)	13
2.3.1 Multiple Kernel Learning (MKL)	15
2.4 Optimasi Hyperparameter	16
2.5 Metrik Evaluasi	18
2.5.1 Accuracy	18
2.5.2 Precision.....	19
2.5.3 Recall	19
2.5.4 F1-Score	19
2.5.5 Area Under the ROC Curve	20
BAB III DESAIN DAN IMPLEMENTASI.....	21
3.1 Desain Penelitian	21
3.1.1 Pengumpulan Data	22
3.1.2 Desain Sistem.....	23
3.1.3 Desain Algoritma	25
3.1.4 Skenario Pengujian	33
3.2 Implementasi.....	38
3.2.1 Implementasi Website	38
3.2.2 Implementasi Algoritma.....	41
BAB IV HASIL DAN PEMBAHASAN.....	46
4.1 Hasil Pengujian	46
4.1.1 Hasil Optimalisasi Hyperparameter	46
4.1.2 Hasil Pengujian K-Fold.....	50
4.1.3 Hasil <i>Feature Importance</i>	52
4.1.4 Hasil Pengujian Subset Fitur	54
4.1.5 Hasil Pengujian Konfigurasi Fitur Tidak Terpilih.....	58

4.1.6 Hasil Perbandingan Kinerja 3 Konfigurasi	61
4.2 Pembahasan.....	64
4.3 Pembahasan Integrasi Islam.....	69
BAB V KESIMPULAN DAN SARAN	73
5.1 Kesimpulan	73
5.2 Saran	74
DAFTAR PUSTAKA	

DAFTAR GAMBAR

Gambar 3. 1 Prosedur Penelitian.....	21
Gambar 3. 2 Desain Sistem.....	24
Gambar 3. 3 Halaman Klasifikasi	39
Gambar 3. 4 Halaman Riwayat Klasifikasi.....	40
Gambar 4. 1 Grafik Algoritma Kneedle Seleksi Fitur.....	53

DAFTAR TABEL

Tabel 2. 1 Penelitian Terkait.....	8
Tabel 3. 1 Struktur Dataset.....	22
Tabel 3. 2 Cuplikan Data.....	23
Tabel 3. 3 Contoh Data Latih dan Data Uji.....	25
Tabel 3. 4 Normalisasi Z-score	27
Tabel 3. 5 Perhitungan Kernel Linear,RBF, dan Gabungan MKL	29
Tabel 3. 6 Fungsi Keputusan Data Uji	32
Tabel 4. 1 Ruang Pencarian Hyperparameter.....	46
Tabel 4. 2 Hasil Grid Search 27 Kombinasi Hyperparameter.....	47
Tabel 4. 3 Hasil Pengujian K-fold.....	50
Tabel 4. 4 Hasil Permutation Feature Importance.....	52
Tabel 4. 5 Hasil Pengujian K-fold Subset Fitur	55
Tabel 4. 6 Hasil Pengujian K-fold Fitur Tidak Terpilih	59
Tabel 4. 7 Perbandingan Kinerja Model pada Tiga Konfigurasi Fitur	61

ABSTRAK

Putra, Muhammad Salman Alfarizy Yonda. 2026. Penerapan Metode **Support Vector Machine dengan Pendekatan Multiple Kernel Learning untuk Klasifikasi Stunting Berbasis Web**. Skripsi. Program Studi Teknik Informatika, Fakultas Sains dan Teknologi, Universitas Islam Negeri Maulana Malik Ibrahim Malang.

Pembimbing: (I) Nurizal Dwi Priandani, M.Kom; (II) Hani Nurhayati, M.T

Kata Kunci: *Support Vector Machine, Multiple Kernel Learning*, stunting, klasifikasi, aplikasi berbasis web

Stunting merupakan kondisi gagal tumbuh pada anak balita akibat kekurangan gizi kronis yang berdampak pada perkembangan fisik dan kognitif secara jangka panjang. Prevalensi stunting di Indonesia masih mencapai 19,8% pada tahun 2024, dan metode deteksi konvensional berbasis pengukuran antropometri manual di Posyandu dinilai rentan terhadap kesalahan subjektif serta lambat dalam pemrosesan data. Penelitian ini menerapkan algoritma *Support Vector Machine* (SVM) dengan pendekatan *Multiple Kernel Learning* (MKL) yang menggabungkan kernel *Linear* dan *Radial Basis Function* (RBF) untuk mengklasifikasikan status stunting pada anak balita secara otomatis dan berbasis web. Data yang digunakan bersumber dari *Indonesian Family Life Survey* gelombang kelima (IFLS-5) dengan sembilan fitur input, meliputi jenis kelamin, umur, tinggi badan, berat badan, pendidikan ibu, status perawatan makanan, status perawatan kesehatan, penerimaan vitamin A, dan imunisasi rotavirus. Optimasi hyperparameter dilakukan menggunakan *grid search* yang dipadukan dengan validasi silang *10-fold*, menghasilkan konfigurasi terbaik pada nilai $C=10$, $\gamma=0,1$, dan $\beta=0,5$. Model dengan sembilan fitur mencapai akurasi sebesar 98,29%, *F1-score* 98,02%, dan AUC 0,998. Analisis *Permutation Feature Importance* melalui algoritma Kneedle mengidentifikasi bahwa hanya tiga fitur antropometri inti, yaitu umur per bulan, tinggi badan, dan jenis kelamin, yang paling dominan, dan model dengan tiga fitur tersebut justru mencapai akurasi 99,3%, *F1-score* 99,4%, serta AUC sempurna 1,000. Model SVM-MKL kemudian diintegrasikan ke dalam aplikasi berbasis web menggunakan Flask dan MySQL untuk mendukung deteksi dini stunting yang mudah diakses oleh tenaga kesehatan di lapangan.

ABSTRACT

Putra, Muhammad Salman Alfarizy Yonda. 2026. **Application of Support Vector Machine Method with Multiple Kernel Learning Approach for Web-Based Stunting Classification**. Thesis. Informatics Engineering Study Program, Faculty of Science and Technology, Universitas Islam Negeri Maulana Malik Ibrahim Malang.

Advisors: (I) Nurizal Dwi Priandani, M.Kom; (II) Hani Nurhayati, M.T

Keywords: *Support Vector Machine, Multiple Kernel Learning*, stunting, classification, web-based application

Stunting is a growth failure condition in children under five caused by chronic malnutrition that has long-term consequences on physical and cognitive development. The prevalence of stunting in Indonesia reached 19.8% in 2024, while conventional detection methods based on manual anthropometric measurements at Posyandu health posts remain prone to subjective errors and slow data processing. This study applies the *Support Vector Machine* (SVM) algorithm with a *Multiple Kernel Learning* (MKL) approach, combining *Linear* and *Radial Basis Function* (RBF) kernels, to automatically classify stunting status in children under five through a web-based system. The data used in this study were sourced from the fifth wave of the *Indonesian Family Life Survey* (IFLS-5), incorporating nine input features comprising gender, age, height, weight, maternal education, food care status, health care status, vitamin A supplementation, and rotavirus immunization. Hyperparameter optimization was conducted using grid search combined with 10-fold cross-validation, yielding the best configuration at $C=10$, $\gamma=0.1$, and $\beta=0.5$. The nine-feature model achieved an accuracy of 98.29%, an F1-score of 98.02%, and an AUC of 0.998. Permutation Feature Importance analysis using the Kneedle algorithm identified that only three core anthropometric features; age in months, height, and gender were the most dominant predictors, and the three-feature model achieved an even higher accuracy of 99.3%, F1-score of 99.4%, and a perfect AUC of 1.000. The SVM-MKL model was subsequently integrated into a web-based application built with Flask and MySQL to support accessible early detection of stunting for health workers in the field.

مستخلص البحث

بوترا، محمد سلمان الفاريزي يوندا. 2026. تطبيق طريقة آلة المتجهات الداعمة بمنهجية التعلم متعدد النوى لتصنيف التقرُّم المعتمد على الويب. رسالة جامعية. برنامج دراسة هندسة المعلوماتية، كلية العلوم والتكنولوجيا، جامعة مولانا مالک إبراهيم الإسلامية الحكومية مالانغ.

المشرفان: نوريزال دوي بريانداني، ماجستير في علوم الحاسوب؛ هاني نورحياتي، ماجستير في الهندسة.

الكلمات المفتاحية: آلة المتجهات الداعمة، التعلم متعدد النوى، التقرُّم، التصنيف، تطبيق قائم على الويب

التقرُّم هو حالة فشل النمو لدى الأطفال دون سن الخامسة الناجمة عن سوء التغذية المزمن، مما يُخلِّف آثاراً بعيدة المدى على النمو الجسدي والمعرفي. بلغت نسبة انتشار التقرُّم في إندونيسيا 19.8% في عام 2024، في حين لا تزال أساليب الكشف التقليدية القائمة على القياسات الأنثروبومترية اليدوية في مراكز بوسيانندو الصحية عُرضةً للأخطاء الذاتية وبطء معالجة البيانات. تُطبَّق هذه الدراسة خوارزمية آلة المتجهات الداعمة (SVM) بمنهجية التعلم متعدد النوى (MKL)، التي تجمع بين النواة الخطية (Linear) ونواة دالة الأساس الشعاعي (RBF)، بهدف تصنيف حالات التقرُّم لدى الأطفال دون سن الخامسة بصورة آلية عبر نظام قائم على الويب. استُقيت البيانات المستخدمة في هذه الدراسة من الموجة الخامسة لمسح الحياة الأسرية الإندونيسي (IFLS-5)، وتشمل تسع سمات مدخلة هي: الجنس، والعمر، والطول، والوزن، ومستوى تعليم الأم، وحالة الرعاية الغذائية، وحالة الرعاية الصحية، وتلقِّي مكملات فيتامين أ، والتطعيم ضد فيروس الروتا. جرى تحسين المعاملات الفائقة باستخدام البحث الشبكي (Grid Search) مقروناً بالتحقق المتقاطع العشري الطيَّات (Fold Cross-Validation-10)، وأسفر ذلك عن أفضل تهيئة عند القيم $C=10$ و $gamma=0.1$ و $beta=0.5$. حقَّق النموذج المؤلَّف من تسع سمات دقةً بلغت 98.29%، ودرجة F1 بلغت 98.02%، ومساحةً تحت منحنى ROC (AUC) بلغت 0.998. كشف تحليل أهمية السمات بالتبديل (Permutation Feature Importance) عبر خوارزمية Kneedle أن ثلاث سمات أنثروبومترية أساسية فحسب وهي العمر بالأشهر، والطول، والجنس هي الأكثر تأثيراً، إذ حقَّق النموذج المقتصر على هذه السمات الثلاث دقةً أعلى بلغت 99.3%، ودرجة F1 بلغت 99.4%، ومساحةً تحت المنحنى مثالية قدرها 1.000. وقد جرى دمج نموذج SVM-MKL في تطبيق قائم على الويب مبنيّ باستخدام Flask وMySQL، لدعم الكشف المبكر عن التقرُّم بصورة يسيرة الوصول لدى العاملين في مجال الصحة في الميدان.

BAB I

PENDAHULUAN

1.1 Latar Belakang

Stunting adalah kondisi gagal tumbuh pada anak yang disebabkan oleh kekurangan gizi kronis. Kondisi ini ditandai dengan tinggi badan yang lebih rendah dibandingkan standar usia, yang dapat menghambat perkembangan fisik, kognitif, dan meningkatkan risiko penyakit kronis (UNICEF, 2021). Di Indonesia, prevalensi stunting pada balita mencapai 19,8% pada tahun 2024, yang menempatkan negara ini sebagai salah satu negara dengan beban stunting tertinggi di Asia Tenggara (Kemenkes RI, 2025; Badan Kebijakan Pembangunan Kesehatan, 2025). Statistik ini menunjukkan urgensi tindakan yang lebih efektif, terutama dalam hal deteksi dan penanganan dini stunting.

Untuk menjawab urgensi tersebut, penanganan stunting memerlukan intervensi komprehensif yang holistik dan berbasis bukti, meliputi pemenuhan gizi adekuat selama masa kehamilan serta dua tahun pertama kehidupan anak, suplementasi mikronutrien, edukasi gizi bagi ibu dan keluarga, dan peningkatan akses pelayanan kesehatan dasar. Dalam perspektif Islam, pendekatan ini selaras dengan kerangka normatif Al-Qur'an yang menekankan tanggung jawab etis orang tua dan masyarakat dalam memastikan penyusuan serta pemeliharaan anak secara optimal sebagai wujud amanah ilahi, sebagaimana firman Allah SWT dalam Surah Al-Baqarah ayat 233:

وَالْوَالِدَتُ يُرْضِعْنَ أَوْلَادَهُنَّ حَوْلَيْنِ كَامِلَيْنِ لِمَنْ أَرَادَ أَنْ يُنِيمَ الرِّضَاعَةُ وَعَلَى الْمَوْلُودِ لَهُ رِزْقُهُنَّ وَكِسْوَتُهُنَّ بِالْمَعْرُوفِ لَا تُكَلَّفُ نَفْسٌ إِلَّا وُسْعَهَا لَا تُضَارَّ وَالِدَةٌ بِوَلَدِهَا وَلَا مَوْلُودٌ لَهُ بِوَلَدِهِ وَعَلَى الْوَارِثِ مِثْلُ ذَلِكَ فَإِنْ أَرَادَا فِصَالًا عَنْ تَرَاضٍ مِنْهُمَا وَتَشَاوُرٍ فَلَا جُنَاحَ عَلَيْهِمَا وَإِنْ أَرَدْتُمْ أَنْ تَسْتَرْضِعُوا أَوْلَادَكُمْ فَلَا جُنَاحَ عَلَيْكُمْ إِذَا سَلَّمْتُمْ مَا آتَيْتُمْ بِالْمَعْرُوفِ وَاتَّقُوا اللَّهَ وَاعْلَمُوا أَنَّ اللَّهَ بِمَا تَعْمَلُونَ بَصِيرٌ

"Dan ibu-ibu (yang menyusukan anaknya) hendaklah menyusukan anak-anaknya selama dua tahun yang penuh, yaitu bagi yang ingin menyempurnakan masa penyusuan. Dan kewajiban ayah memberi makan dan pakaian kepada para ibu (yang menyusukan anaknya) menurut yang patut. Seseorang tidak diwajibkan atas sesuatu yang melebihi kesanggupannya. Janganlah seorang ibu menderita karena anaknya dan seorang ayah (juga) karena anaknya, dan kewajiban waris pun demikian. Jika keduanya ingin menyapih (sebelum dua tahun) dengan musyawarah dan persetujuan keduanya, maka tidak ada dosa baginya. Dan jika kamu ingin anak-anakmu disusukan oleh orang lain, tidak ada dosa bagimu, jika kamu menunaikan pembayaran dengan cara yang patut. Dan bertakwalah kepada Allah, dan ketahuilah bahwa Allah Maha Melihat apa yang kamu kerjakan." (QS. Al-Baqarah 2:233)

Ayat ini menegaskan secara jelas kewajiban orang tua untuk menyusui anak secara optimal selama dua tahun penuh, dengan berlandaskan pada prinsip keadilan (ma'ruf) serta tanggung jawab kolektif dalam pemenuhan kebutuhan nutrisi dan perawatan, sebagaimana diamanatkan melalui perintah takwa kepada Allah SWT. M. Quraish Shihab dalam Tafsir Al-Mishbah (Jilid 1, hlm. 503-506) menjelaskan bahwa ayat ini menggunakan redaksi berita (khabar) yang berfungsi sebagai perintah yang sangat kuat. Kata *al-wālidāt* (para ibu) tidak terbatas pada ibu kandung saja, tetapi mencakup pula ibu susuan. Namun demikian, air susu ibu kandung tetap lebih utama karena bayi telah mengenal detak jantung ibunya sejak dalam kandungan. Ayat ini juga menegaskan bahwa masa menyusui yang sempurna adalah dua tahun penuh, yang merupakan batas maksimal. Kewajiban memberi makan dan pakaian kepada ibu yang menyusui berada di pundak ayah, sesuai

dengan kemampuannya. Quraish Shihab menggarisbawahi bahwa penyusuan bukan hanya persoalan pemenuhan gizi jasmani, tetapi juga pendidikan ruhani yang fundamental bagi pembentukan karakter anak. Dalam perspektif Islam, ketentuan ini menjadi dasar etis dalam upaya pencegahan stunting, di mana pemenuhan gizi sejak dini berperan penting dalam mencegah kekurangan nutrisi kronis yang dapat menghambat pertumbuhan anak. Hal ini sejalan dengan bukti ilmiah yang menunjukkan bahwa periode awal kehidupan merupakan fase kritis perkembangan (UNICEF, 2021).

Meskipun demikian, implementasi deteksi dini stunting melalui model SVM-MKL berbasis web seperti yang diusulkan dalam penelitian ini masih dihadapkan pada tantangan metode konvensional, yaitu pengukuran antropometri manual di Posyandu, yang rentan terhadap kesalahan subjektif dan keterlambatan pemrosesan data, khususnya di wilayah terpencil dengan infrastruktur terbatas (Kemenkes RI, 2022). Akibatnya, pendekatan berbasis algoritma pembelajaran mesin menjanjikan peningkatan efisiensi dan akurasi, sebagaimana dibuktikan dalam studi kasus medis heterogen (Rahmi et al., 2022). Di sisi lain, upaya pemerintah Indonesia melalui program pemantauan pertumbuhan anak, edukasi gizi, dan suplementasi makanan tambahan telah menunjukkan kemajuan, namun tetap terhambat oleh kekurangan tenaga ahli terlatih serta ketidakmerataan distribusi layanan, yang menghalangi pencapaian target penurunan prevalensi stunting di bawah 14,2% pada tahun 2029 sesuai Rencana Pembangunan Jangka Menengah Nasional (Kementerian PPN, 2025). Oleh karenanya, pengembangan

inovasi teknologi yang skalabel dan aksesibel menjadi imperatif untuk mengoptimalkan deteksi serta intervensi stunting secara berkelanjutan.

Di antara berbagai inovasi teknologi tersebut, salah satu pendekatan yang potensial adalah algoritma *Support Vector Machine* (SVM), metode pembelajaran mesin yang melakukan klasifikasi data dengan mencari hyperplane pemisah optimal berdasarkan pola fitur, sebagaimana pertama kali diperkenalkan oleh (Cortes & Vapnik, 1995). Dalam aplikasi stunting, SVM dapat menganalisis faktor risiko seperti usia, berat lahir, dan lingkaran lengan atas untuk memprediksi status stunting anak secara cepat dan akurat, dengan tingkat akurasi mencapai 85-90% pada dataset medis heterogen (Rahmi et al., 2022). Pendekatan ini terbukti efektif bagi tenaga kesehatan dalam pengambilan keputusan pencegahan, terutama ketika terintegrasi dengan data survei nasional, sehingga mengatasi keterbatasan aksesibilitas di daerah terpencil yang disebutkan sebelumnya.

Pendekatan Multiple Kernel Learning (MKL) dikembangkan sebagai solusi terhadap keterbatasan Support Vector Machine (SVM) tunggal dalam mengolah data heterogen yang kompleks. Secara konseptual, MKL memungkinkan kombinasi beberapa fungsi kernel yang dioptimalkan berdasarkan bobot kontribusi masing-masing terhadap peningkatan performa klasifikasi (Gönen & Alpaydm, 2011). Fleksibilitas ini memberikan keunggulan dalam pemrosesan data multidimensi, seperti penggunaan kernel linier untuk fitur numerik dan kernel sparse untuk data kategorikal, yang terbukti dapat meningkatkan akurasi prediksi SVM hingga 5–10% pada kasus medis, termasuk klasifikasi gizi (Mariette & Villa-Vialaneix, 2018). Hasil penelitian terdahulu oleh Rahmi et al. (2022) dan Adzhima (2023)

menunjukkan potensi SVM dalam mengklasifikasikan status stunting anak dengan tingkat akurasi yang tinggi. Namun, pendekatan berbasis single kernel masih memiliki keterbatasan dalam generalisasi terhadap fitur campuran serta adaptivitas terhadap pola non-linier. Berdasarkan hal tersebut, penelitian ini mengusulkan penerapan model SVM-MKL dengan kombinasi kernel Linier dan RBF melalui optimasi bobot untuk meningkatkan robustness dan akurasi klasifikasi. Model ini diintegrasikan ke dalam sistem berbasis web agar dapat mendukung deteksi dini stunting secara lebih efektif dan mudah diakses.

Dengan demikian, melalui penelitian ini diharapkan terwujud model SVM-MKL yang tidak hanya memvalidasi efektivitasnya dalam prediksi stunting dengan akurasi di atas 95% pada dataset IFLS-5 yang lebih luas dan beragam, tetapi juga menghadirkan solusi praktis berbasis web yang siap diadopsi oleh Posyandu dan orang tua, sehingga meningkatkan efisiensi deteksi secara keseluruhan serta mendukung target nasional penurunan prevalensi stunting di bawah 14,2% pada 2029 (Kementerian PPN, 2025).

1.2 Rumusan Masalah

Berdasarkan latar belakang yang telah diuraikan, rumusan masalah penelitian ini dapat dirumuskan sebagai berikut:

1. Bagaimana implementasi model klasifikasi status stunting pada anak balita menggunakan algoritma *Support Vector Machine* (SVM) dengan pendekatan *Multiple Kernel Learning* (MKL) yang terintegrasi ke dalam aplikasi berbasis web?

2. Bagaimana kinerja algoritma *Support Vector Machine* (SVM) dengan pendekatan *Multiple Kernel Learning* (MKL) pada kasus stunting?

1.3 Tujuan Penelitian

Adapun tujuan penelitian ini mencakup hal-hal sebagai berikut:

1. Mengimplementasikan model klasifikasi status stunting pada anak balita menggunakan algoritma *Support Vector Machine* (SVM) dengan pendekatan *Multiple Kernel Learning* (MKL), yang terintegrasi ke dalam aplikasi berbasis web untuk memfasilitasi penggunaan praktis dalam penilaian stunting.
2. Mengevaluasi kinerja algoritma *Support Vector Machine* (SVM) dengan pendekatan *Multiple Kernel Learning* (MKL) dalam klasifikasi status stunting pada anak balita, melalui metrik evaluasi seperti akurasi, presisi, recall, F1-score, guna menentukan efektivitas dan keunggulannya dalam menangani kasus stunting.

1.4 Manfaat Penelitian

Penelitian ini diharapkan dapat memberikan manfaat sebagai berikut:

1. Memberikan kontribusi pada pengembangan metode pembelajaran mesin untuk klasifikasi data medis heterogen, sebagai referensi bagi penelitian di Teknik Informatika dan kesehatan.
2. Menyediakan alat deteksi dini stunting berbasis web yang mudah digunakan oleh tenaga kesehatan, untuk mendukung intervensi gizi anak secara cepat dan akurat.

3. Mendorong pencegahan stunting di masyarakat melalui teknologi aksesibel, serta menjadi bahan ajar untuk pendidikan interdisipliner di bidang informatika dan kesehatan masyarakat.

1.5 Batasan Penelitian

Penelitian ini menetapkan batasan-batasan tertentu guna memfokuskan analisis secara lebih terarah. Adapun batasan-batasan masalah yang diterapkan dalam penelitian ini adalah sebagai berikut:

1. Data penelitian ini menggunakan Indonesian Family Life Survey gelombang kelima (IFLS-5) tahun 2014–2015 yang diselenggarakan oleh RAND.
2. Status stunting dikategorikan menjadi 2 kelas, yaitu stunting dan Normal.

BAB II

STUDI PUSTAKA

2.1 Penelitian Terkait

Kajian terhadap berbagai studi yang relevan telah dilakukan guna memperkuat landasan teori serta mendukung hasil yang diharapkan. Tabel 2.1 menyajikan daftar penelitian yang memiliki kesamaan dalam penggunaan metode *Support Vector Machine* (SVM) untuk klasifikasi stunting. Setiap penelitian menerapkan SVM dengan jenis kernel yang berbeda. Fokus utama dari penelitian ini adalah penerapan SVM menggunakan pendekatan *Multiple Kernel Learning* (MKL), yang melibatkan penggunaan kernel linear dan kernel Radial Basis Function (RBF) dalam menyelesaikan permasalahan klasifikasi stunting.

Tabel 2. 1 Penelitian Terkait

No	Peneliti	Metode	Variable/Input	Hasil (angka)
1	(Angga Aditya Permana, 2024)	SVM kernel linear	Umur, berat badan, tinggi badan, dll	Akurasi 82%, Presisi 80%, Recall 86%
2	(Gaffar et al., 2024)	SVM kernel polynomial	Data balita (umur, TB, BB, dll)	Akurasi 98% (rata-rata 96,98%), Presisi 96,99%, Recall 96,98%, F1 96,94%
3	(Rahmi et al., 2022)	SVM kernel RBF (C=10, gamma=5)	Faktor risiko stunting (umur, TB, BB, jenis kelamin)	Akurasi 100%
4	(Hasdyna et al., 2024)	SVM kernel RBF dan Sigmoid	Data stunting Aceh	RBF 91,3%, Sigmoid 85,6%
5	(Adzhima, 2023)	SVM kernel polynomial (gamma=0,01)	Data balita	Akurasi rata-rata 98,99%, akurasi terbaik 100%
6	(Muhammad Athoillah, 2018)	Multi Kernel SVM (linear, polynomial, RBF)	Fitur citra kendaraan bermotor	Akurasi 84,60%, Presisi 84,95%, Recall 84,60%
7	(Yunita et al., 2012)	Multiple Kernel SVM	Data ekspresi gen (leukimia & tumor usus besar)	Akurasi 85% (leukimia), 69% (tumor usus besar)

Penelitian oleh (Angga Aditya Permana, 2024) berfokus pada penggunaan algoritma *Support Vector Machine* (SVM) untuk mengklasifikasikan status stunting pada balita. Penelitian ini menangani permasalahan klasifikasi yang dilakukan secara manual, yang memakan waktu dan rentan kesalahan. Dengan menggunakan SVM kernel linear, penelitian ini berhasil mencapai akurasi 82%, presisi 80%, dan recall 86%. Model klasifikasi ini diintegrasikan ke dalam aplikasi berbasis web menggunakan Streamlit, yang membuatnya lebih mudah diakses untuk deteksi stunting secara real-time.

Selanjutnya, penelitian oleh (Gaffar et al., 2024) menggunakan SVM dengan kernel polynomial untuk klasifikasi stunting di Kabupaten Enrekang. Penelitian ini memperoleh hasil yang sangat baik dengan akurasi tertinggi 99,13% pada fold ke-4 dan nilai rata-rata akurasi 96,98%. Hasil ini menunjukkan efektivitas SVM dengan kernel polynomial dalam mengklasifikasikan stunting dengan akurat. Temuan ini mendukung penggunaan SVM untuk analisis stunting yang lebih efisien.

Selain itu, penelitian oleh (Rahmi et al., 2022) menerapkan SVM dengan kernel *Radial Basis Function* (RBF) untuk mengklasifikasikan stunting pada balita di Padang. Penelitian ini menunjukkan bahwa model SVM dengan kernel RBF dapat mencapai akurasi 100% dalam memprediksi stunting. Hal ini mengindikasikan bahwa SVM sangat efektif untuk deteksi dini masalah stunting pada balita. Penggunaan pembelajaran mesin dalam bidang kesehatan terbukti dapat meningkatkan akurasi prediksi.

Selanjutnya, dalam penelitian oleh (Hasdyna et al., 2024), pendekatan hibrid pembelajaran mesin digunakan untuk klasifikasi, prediksi, dan optimasi clustering prevalensi stunting di Aceh. SVM dengan kernel RBF menunjukkan akurasi 91,3%, lebih tinggi dibandingkan dengan kernel sigmoid yang hanya mencapai 85,6%. Penelitian ini juga mengoptimalkan clustering dengan metode K-Medoids, meningkatkan pemahaman tentang distribusi stunting. Hasil ini berkontribusi pada perencanaan intervensi kesehatan masyarakat yang lebih terarah.

Berikutnya, Adzhima (2023) mengembangkan sistem klasifikasi status stunting balita berbasis web dengan menggunakan algoritma SVM sebagai mesin utama klasifikasi. Dengan pemilihan kernel polynomial dan pengaturan parameter $\gamma = 0,01$, penelitian ini memperoleh akurasi rata-rata 98,99% dengan performa terbaik mencapai 100%, serta mengimplementasikan model dalam aplikasi web untuk mempermudah deteksi dini stunting.

Selanjutnya, Athoillah (2018) menerapkan metode Multi Kernel Support Vector Machine untuk mengklasifikasikan kendaraan bermotor roda dua dan roda empat pada gerbang tol otomatis, sebagai solusi atas kebutuhan sistem yang mampu membedakan jenis kendaraan secara otomatis. Dengan mengombinasikan kernel linear, polynomial orde tiga, dan RBF dalam kerangka Multiple Kernel Learning, model klasifikasi yang dibangun mampu mencapai akurasi rata-rata 84,60%, presisi 84,95%, dan recall 84,60%, sehingga menunjukkan efektivitas pendekatan multi kernel pada data citra nonlinier.

Selain itu, (Yunita et al., 2012) mengimplementasikan Multiple Kernel Support Vector Machine untuk seleksi fitur pada data ekspresi gen dengan studi

kasus leukemia dan tumor usus besar, guna mengatasi masalah jumlah fitur yang sangat besar dengan jumlah sampel yang relatif sedikit. Pendekatan ini menggunakan kombinasi beberapa kernel untuk mengeliminasi fitur yang kurang relevan, menghasilkan akurasi 85% pada data leukemia dan 69% pada data tumor usus besar, yang lebih tinggi dibanding metode SVM dasar dengan akurasi 82% dan 59%, serta sekaligus mempercepat waktu komputasi pada proses klasifikasi.

Berdasarkan penelitian terdahulu, sebagian besar studi klasifikasi stunting masih menggunakan SVM dengan single kernel, seperti *linear*, *polynomial*, atau *RBF*, sedangkan pemanfaatan pendekatan *Multiple Kernel Learning* pada kasus stunting masih terbatas. Oleh karena itu, penelitian ini menawarkan kontribusi dalam bentuk penerapan kombinasi kernel *linear* dan *RBF* pada *SVM-MKL*, analisis perbandingan variasi set fitur, serta implementasi model ke dalam aplikasi berbasis website.

2.2 Stunting

Stunting merupakan kondisi gagal tumbuh pada anak balita yang diakibatkan oleh kekurangan gizi kronis (*chronic malnutrition*), yaitu defisiensi berkelanjutan zat gizi makro dan mikro selama periode kritis 1.000 hari pertama kehidupan (konsepsi hingga usia dua tahun), yang secara antropometri ditandai dengan Height-for-Age Z-score (HAZ) < -2 SD dari mediana standar WHO berdasarkan usia dan jenis kelamin. Kekurangan gizi kronis bersifat kumulatif dan irreversibel, membedakannya dari malnutrisi akut (*wasting*) yang bersifat sementara (Anggraini & Rachmawati, 2021). Prevalensi stunting nasional Indonesia mencapai 19,8% pada tahun 2024, meskipun angka ini masih di bawah

proyeksi Bappenas sebesar 20,1% untuk periode yang sama (Kemenkes RI, 2025; Badan Kebijakan Pembangunan Kesehatan, 2025)..

Penyebab langsung stunting terdiri dari asupan gizi tidak adekuat secara kronis dan infeksi berulang yang mengganggu penyerapan nutrisi, sementara faktor tidak langsung meliputi pendidikan ibu rendah, sanitasi buruk, kemiskinan keluarga, serta akses layanan kesehatan terbatas (Hidayati et al., 2022). Faktor risiko signifikan mencakup kurangnya pemahaman ibu tentang pola gizi seimbang, yang menyebabkan ketidakmampuan memenuhi kebutuhan nutrisi anak secara optimal (Sugiarti et al., 2023). Indikator antropometri primer stunting adalah TB/U (HAZ), yang didukung oleh pengukuran BB/U (*underweight*) dan BB/TB (*wasting*) untuk diagnosis status stunting komprehensif.

Penurunan angka stunting memerlukan pendekatan holistik yang melibatkan edukasi gizi ibu, pemberian makanan tambahan (PMT) untuk balita, peningkatan sanitasi, dan akses layanan kesehatan dasar. Intervensi berbasis masyarakat seperti *positive deviance* terbukti efektif mengidentifikasi praktik pemberian makan dan pengasuhan sukses di lingkungan berpenghasilan rendah (Sugiarti et al., 2023). Namun demikian, tantangan utama di Indonesia tetap pada ketidakmerataan distribusi layanan di wilayah terpencil (Fardila elba et al., 2023). Oleh karena itu, pengembangan sistem deteksi dini berbasis teknologi seperti aplikasi berbasis web *SVM MKL* menjadi imperatif untuk mempercepat identifikasi risiko stunting dan memfasilitasi intervensi tepat sasaran guna mendukung target penurunan prevalensi nasional (Kementerian PPN, 2025).

2.3 Support Vector Machine (SVM)

Support Vector Machine (SVM) adalah algoritma pembelajaran mesin yang banyak digunakan dalam masalah klasifikasi dan regresi. SVM bekerja dengan prinsip *Structural Risk Minimization* (SRM) untuk mengoptimalkan margin pemisah antar kelas data. Tujuan utama dari SVM adalah untuk menemukan hyperplane terbaik yang memisahkan data ke dalam dua kelas dengan margin terbesar. Dalam konsep ini, hyperplane adalah sebuah bidang pemisah dalam ruang fitur yang membagi data ke dalam dua kelompok berbeda (Cortes & Vapnik, 1995). SVM berusaha untuk memaksimalkan margin antara dua kelas yang ada. Margin adalah jarak terpendek antara hyperplane dan titik data terdekat dari masing-masing kelas. Titik data yang berada pada batas margin ini disebut sebagai *support vectors*, yang memiliki peran penting dalam menentukan posisi *hyperplane*.

Fungsi keputusan pada SVM adalah rumus utama yang digunakan untuk memprediksi kelas dari data uji berdasarkan nilai kernel antara data latih dan data uji. Fungsi keputusan ini dapat ditulis sebagai berikut:

$$f(x) = \sum_{i=1}^n \alpha_i y_i K(x_i, x) + b \quad (2.1)$$

Dimana $f(x)$ merupakan hasil prediksi fungsi keputusan untuk data uji x , α_i adalah *Lagrange multipliers* yang dihasilkan dari proses pelatihan SVM, yang menunjukkan seberapa penting setiap sampel dalam menentukan pemisahan kelas, y_i adalah label kelas untuk data pelatihan x_i , yang dapat bernilai -1 atau +1, $K(x_i, x)$ adalah fungsi kernel yang mengukur kesamaan antara data pelatihan x_i dan data uji x , memungkinkan SVM untuk menangani data yang tidak dapat dipisahkan

secara linier, b adalah bias atau intercept yang menentukan posisi hyperplane dalam ruang fitur, n adalah jumlah sampel dalam dataset pelatihan. Rumus ini menjelaskan bahwa fungsi keputusan pada SVM bergantung pada *support vectors* yang memiliki bobot α_i dan juga nilai dari kernel yang digunakan untuk menghitung kesamaan antar titik data (Cortes & Vapnik, 1995).

Salah satu aspek penting dari SVM adalah penggunaan kernel, yang memungkinkan SVM untuk mengatasi masalah data yang tidak dapat dipisahkan secara linier. Kernel digunakan untuk menghitung kesamaan antar titik data dalam ruang fitur yang lebih tinggi, di mana pemisahan kelas dapat dilakukan dengan lebih mudah. Terdapat beberapa jenis kernel yang sering digunakan dalam SVM, di antaranya kernel linear dan kernel Radial Basis Function (RBF). Kernel linear adalah jenis kernel yang paling sederhana dalam konteks SVM. Fungsi kernel linear dapat dinyatakan sebagai berikut:

$$K_{linear}(x_i, x_j) = x_i^T x_j \quad (2.2)$$

Dimana x_i dan x_j adalah vektor fitur dua titik data dalam ruang input. $x_i^T x_j$ adalah hasil perkalian dot antara dua vektor fitur. Hasil ini mengukur kedekatan atau kesamaan antara dua titik data. Kernel linear sangat efektif pada dataset yang terpisah secara linier, karena tidak memerlukan transformasi data lebih lanjut. Kernel ini cukup sederhana dan efisien dalam hal komputasi, serta mudah diinterpretasikan. Namun, kernel ini terbatas pada data yang dapat dipisahkan secara linier (Schölkopf & Smola, 2001).

Di sisi lain, kernel RBF (Radial Basis Function) digunakan ketika data tidak terpisah secara linier. Fungsi kernel RBF dapat didefinisikan sebagai berikut:

$$K_{RBF}(x_i, x_j) = \exp(-\gamma \|x_i - x_j\|^2) \quad (2.3)$$

Dimana $\|x_i - x_j\|^2$ adalah jarak Euclidean antara dua titik data x_i dan x_j . γ adalah parameter yang mengontrol lebar fungsi Gaussian. Semakin besar nilai γ , semakin sempit pengaruh tiap titik data. $\exp(-\gamma \|x_i - x_j\|^2)$ adalah fungsi eksponensial yang mengubah data ke dalam dimensi lebih tinggi, sehingga mempermudah pemisahan non-linier.

Kernel RBF sangat berguna dalam menghadapi dataset yang memiliki pola non-linier, karena ia memungkinkan pemisahan yang lebih fleksibel antara kelas-kelas data. Namun, kernel RBF membutuhkan penyetelan parameter γ yang tepat agar tidak terjadi *overfitting* atau *underfitting* pada model (Schölkopf & Smola, 2001).

2.3.1 Multiple Kernel Learning (MKL)

Data yang memiliki karakteristik yang sangat beragam sering kali tidak dapat dimodelkan dengan baik hanya oleh satu jenis kernel. Multiple Kernel Learning (MKL) hadir sebagai pendekatan yang menggabungkan beberapa fungsi kernel untuk meningkatkan kinerja SVM. Alih-alih memilih satu kernel terbaik, MKL menggunakan beberapa kernel sekaligus dan memberikan bobot yang menunjukkan kontribusi masing-masing kernel terhadap fungsi keputusan akhir.

Fungsi kernel gabungan dalam MKL ditulis sebagai:

$$K_{MKL}(x_i, x_j) = \beta \cdot K_{\text{linear}}(x_i, x_j) + (1 - \beta) \cdot K_{RBF}(x_i, x_j) \quad (2.4)$$

Persamaan (2.4) mendefinisikan kernel gabungan MKL sebagai kombinasi linear kernel linear dan RBF. Bobot β ($0 \leq \beta \leq 1$) mengatur kontribusi kernel linear, sementara $(1-\beta)$ mengatur kontribusi kernel RBF. Nilai ekstrem $\beta=1$ atau $\beta=0$ masing-masing hanya mengaktifkan kernel linear atau RBF, sedangkan nilai di antara keduanya mengaktifkan kombinasi keduanya. Dalam kerangka ini, kernel linear memodelkan hubungan linear antarfitur, dan kernel RBF memodelkan hubungan non-linear yang kompleks. Kombinasi ini unggul pada data heterogen yang mengandung kedua pola hubungan tersebut.

Bobot β dioptimalkan selama proses pelatihan bersamaan dengan parameter C dan γ menggunakan teknik *grid search* yang dipadukan dengan validasi silang. Kernel linear bertugas menangani hubungan-hubungan linear antarfitur, sementara kernel RBF menangkap pola-pola non-linear yang lebih rumit. Kombinasi optimal dari kedua kernel ini memungkinkan SVM memperoleh hasil yang lebih baik dibandingkan dengan penggunaan kernel tunggal pada data yang kompleks dan heterogen (Gönen & Alpaydin, 2011).

2.4 Optimasi Hyperparameter

Optimasi *hyperparameter* adalah langkah penting dalam pembelajaran mesin yang bertujuan menemukan kombinasi nilai parameter terbaik sebelum proses pelatihan dimulai. Parameter seperti C , γ , dan β tidak dipelajari secara otomatis dari data, melainkan harus ditetapkan terlebih dahulu. Pemilihan nilai yang tepat sangat menentukan kemampuan model dalam menangkap pola tanpa kehilangan daya generalisasi, karena *hyperparameter* mengendalikan keseimbangan antara bias dan varians (Andi et al., 2026).

Salah satu teknik yang paling sistematis untuk melakukan optimasi ini adalah *grid search*. Teknik ini membentuk kisi-kisi dari nilai-nilai *hyperparameter* yang ingin diuji, lalu mencoba setiap kombinasi yang mungkin. Setiap kombinasi dievaluasi menggunakan validasi silang, dan kombinasi yang menghasilkan metrik evaluasi tertinggi dipilih sebagai konfigurasi optimal. *Grid search* menjamin bahwa seluruh kemungkinan dalam ruang pencarian telah dijelajahi. Konsekuensinya, waktu komputasi dapat bertambah besar seiring bertambahnya jumlah *hyperparameter* dan variasi nilai yang diuji (Nghien et al., 2025).

Penelitian yang dilakukan oleh (Nghien et al., 2025) membandingkan kinerja *grid search* dengan empat algoritma optimasi lain, yaitu *random search*, *Bayesian optimization*, *genetic algorithm*, dan *chemical reaction optimization*, pada optimasi SVM. Hasil pengujian mereka menunjukkan bahwa *grid search* tetap mampu menghasilkan akurasi yang bersaing, meskipun dari sisi waktu bukanlah yang tercepat.

Penelitian ini mengoptimasi tiga *hyperparameter* utama SVM-MKL melalui *grid search*. Ketiga parameter tersebut adalah C (koefisien regularisasi), γ (lebar fungsi Gaussian pada kernel RBF), dan β (bobot kontribusi kernel linear). Parameter C mengatur seberapa besar toleransi model terhadap kesalahan klasifikasi. Parameter γ menentukan seberapa luas satu titik data memengaruhi permukaan keputusan. Sementara itu, β mengontrol keseimbangan peran kernel linear dan kernel RBF di dalam MKL. Optimasi

dilakukan secara serentak agar ditemukan satu kombinasi yang memberikan performa klasifikasi stunting tertinggi.

2.5 Metrik Evaluasi

Kinerja model klasifikasi perlu diukur menggunakan metrik yang mampu menilai berbagai aspek secara objektif. Penelitian ini memakai lima metrik, yaitu accuracy, precision, recall, F1-score, dan Area Under the ROC Curve (AUC). Masing-masing metrik menyoroti sisi yang berbeda dari kualitas prediksi, sehingga jika digunakan bersama-sama akan menghasilkan penilaian yang menyeluruh.

2.5.1 Accuracy

Accuracy mengukur persentase prediksi yang benar dari seluruh sampel yang diuji. Perhitungannya didasarkan pada empat komponen confusion matrix:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2.5)$$

True Positive (TP) adalah jumlah balita stunting yang tepat diprediksi stunting. True Negative (TN) adalah jumlah balita normal yang tepat diprediksi normal. False Positive (FP) adalah jumlah balita normal yang keliru diprediksi stunting. False Negative (FN) adalah jumlah balita stunting yang keliru diprediksi normal. Accuracy memberikan ringkasan yang sederhana, tetapi dapat menyesatkan jika proporsi kedua kelas sangat timpang (Chicco & Jurman, 2020).

2.5.2 Precision

Precision mengukur seberapa tepat model saat menyatakan seorang balita mengalami stunting. Rumusnya adalah:

$$Precision = \frac{TP}{TP+FP} \quad (2.6)$$

Nilai precision yang tinggi berarti model jarang menghasilkan false positive. Dalam deteksi stunting, keadaan ini membuat orang tua tidak menerima peringatan yang keliru, dan tenaga kesehatan tidak membuang sumber daya untuk menangani anak yang sebenarnya normal (Hicks et al., 2022).

2.5.3 Recall

Recall atau sensitivity mengukur seberapa lengkap model dalam menemukan kasus stunting yang sesungguhnya. Rumusnya adalah:

$$Recall = \frac{TP}{TP+FN} \quad (2.7)$$

Nilai recall yang tinggi berarti sangat sedikit kasus stunting yang terlewatkan (false negative). Setiap false negative berarti seorang anak yang membutuhkan bantuan tidak terdeteksi (Hicks et al., 2022).

2.5.4 F1-Score

F1-Score menyatukan precision dan recall ke dalam satu angka melalui rata-rata harmonis. Rata-rata harmonis dipilih karena ia memberi penalti lebih besar jika salah satu dari precision atau recall bernilai rendah. Dengan demikian, F1-Score hanya akan tinggi jika keduanya sama-sama tinggi. Rumusnya adalah:

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.8)$$

F1-Score sangat berguna untuk menilai model pada data dengan kelas yang tidak seimbang (Hicks et al., 2022). Dalam penelitian ini, F1-Score dijadikan acuan utama karena mencerminkan keseimbangan antara risiko false positive dan false negative di lapangan.

2.5.5 Area Under the ROC Curve

Area Under the ROC Curve (AUC) mengukur kemampuan model membedakan kelas stunting dan normal tanpa terikat pada satu ambang batas tertentu. Nilai AUC 1 berarti diskriminasi sempurna, sedangkan 0,5 setara dengan tebakan acak (Bradley, 1997). Keunggulan AUC terletak pada sifatnya yang tidak terpengaruh oleh biaya relatif kesalahan klasifikasi, memiliki kepekaan tinggi dalam membandingkan model, dan tidak bergantung pada proporsi kelas dalam data.

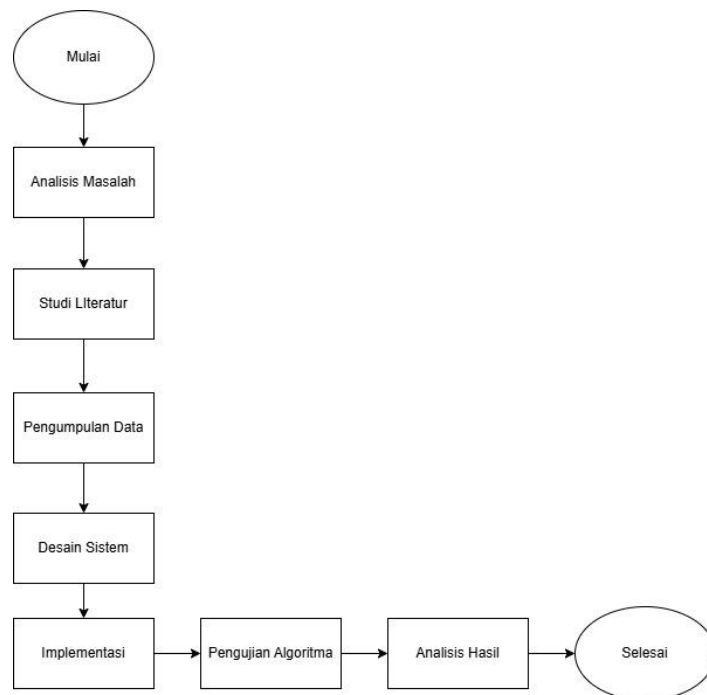
Fawcett (2006) menegaskan bahwa ROC dan AUC telah menjadi alat baku evaluasi model klasifikasi karena mampu merangkum performa di semua kemungkinan ambang batas dalam satu angka. Fleksibilitas ini memberi keleluasaan bagi petugas kesehatan untuk menyesuaikan titik potong klasifikasi sesuai dengan kondisi sumber daya di lapangan.

BAB III

DESAIN DAN IMPLEMENTASI

3.1 Desain Penelitian

Pada tahap ini, disusun desain penelitian yang akan menjadi pedoman dalam pelaksanaan setiap tahapan penelitian, yang mencakup analisis masalah, kajian pustaka, pengumpulan data, perancangan sistem, implementasi sistem, pengujian algoritma, serta analisis hasil. Penyusunan desain ini bertujuan untuk memastikan bahwa setiap tahapan penelitian dapat dilaksanakan dengan terstruktur dan sistematis, sekaligus mempermudah dalam pemantauan dan evaluasi proses penelitian. Gambar 3.1 menyajikan representasi grafis dari alur penelitian yang menggambarkan tahapan-tahapan yang dimulai dari tahap awal hingga mencapai tahap akhir.



Gambar 3. 1 Prosedur Penelitian

3.1.1 Pengumpulan Data

Penelitian ini menggunakan data sekunder yang bersumber dari Indonesian Family Life Survey gelombang kelima (IFLS-5) yang dikumpulkan pada periode 2014–2015. IFLS-5 dipilih karena memiliki cakupan nasional dengan desain multistage stratified sampling serta menyediakan informasi rumah tangga dan individu yang cukup rinci untuk menganalisis faktor-faktor yang berhubungan dengan stunting pada anak.

Variabel yang digunakan dalam penelitian ini mencakup identitas teknis dan keterkaitan anak dengan rumah tangga, karakteristik demografis, indikator antropometri, serta variabel pengasuhan gizi dan kesehatan. Identitas teknis direpresentasikan oleh kolom `child_id` yang berfungsi untuk mengidentifikasi setiap anak secara unik sekaligus menghubungkannya dengan unit rumah tangga asal. Karakteristik demografis anak dideskripsikan melalui umur per bulan untuk memastikan kesesuaian usia saat pengukuran. Kondisi antropometri anak digambarkan oleh tinggi badan dan berat badan yang menjadi dasar penentuan status stunting. Sementara itu, konteks pengasuhan gizi dan kesehatan tercermin pada variabel pendidikan ibu, status perawatan makanan anak, status perawatan kesehatan anak, penerimaan vitamin A, dan imunisasi rotavirus.

Tabel 3. 1 Struktur Dataset

No	Kolom	Keterangan
1.	<code>child_id</code>	ID unik untuk setiap anak dalam dataset.
2.	Jenis Kelamin	Jenis kelamin anak (Laki-laki/Perempuan).
3.	Umur Per Bulan	Usia anak dalam satuan bulan.
4.	Tinggi	Tinggi badan anak (cm).
5.	Berat	Berat badan anak (kg).
6.	Pendidikan Ibu	Tingkat pendidikan formal ibu anak.
7.	Status Perawatan Makanan anak	Status/kualitas perawatan terkait pemberian makan dan asupan gizi anak.

No	Kolom	Keterangan
8.	Status Perawatan Kesehatan Anak	Status/kualitas perawatan kesehatan yang diterima anak.
9.	Menerima vitamin A	Indikator berapa kali anak menerima suplementasi vitamin A.
10.	Imunisasi Rotavirus 1	Indikator apakah anak mendapatkan imunisasi rotavirus (ya/tidak).

Struktur di atas menunjukkan bahwa setiap baris merepresentasikan satu anak yang dikaitkan dengan rumah tangga tertentu dan memiliki informasi lengkap mengenai identitas, usia, kondisi antropometri, serta konteks pengasuhan gizi dan kesehatan.

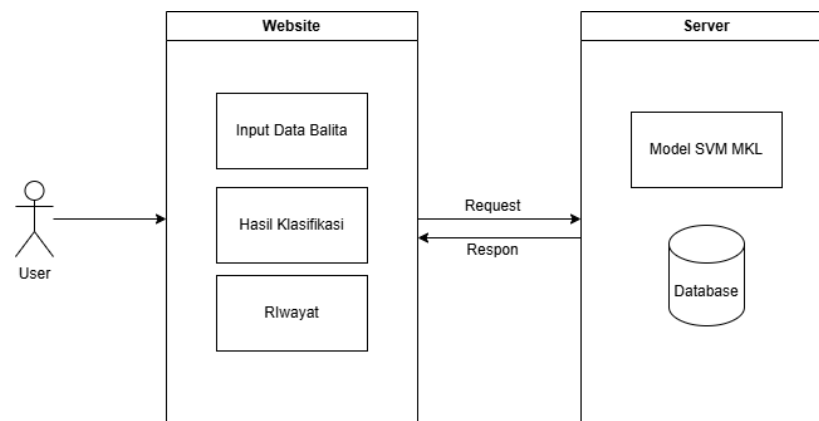
Tabel 3. 2 Cuplikan Data

child_id	Jenis Kelamin	Umur Per Bulan	Tinggi (cm)	Berat (kg)	Pendidikan Ibu	Status Perawatan Makanan	Status Perawatan Kesehatan	Menerima vitamin A	Imunisasi Rotavirus 1
1001	1	36	92	13.0	60	3	3	3	1
1002	1	30	84	10.5	61	2	2	0	3
1003	1	36	86	11	60	3	3	3	1

Contoh data pada tabel 3.2 memberikan gambaran konkret mengenai bentuk observasi yang diolah, mulai dari identitas teknis, profil demografis, kondisi antropometri, hingga status pengasuhan dan label stunting. Nilai umur, tinggi, dan berat badan menjadi dasar penentuan status stunting, sedangkan variabel lain digunakan untuk menangkap variasi faktor risiko yang kemudian dieksplorasi oleh model SVM-MKL dalam proses pelatihan dan pengujian.

3.1.2 Desain Sistem

Berikut merupakan desain sistem yang digunakan dalam penelitian ini seperti diagram blok pada gambar 3.2:



Gambar 3. 2 Desain Sistem

Ilustrasi pada Gambar 3.2 menunjukkan alur kerja sistem yang dirancang dalam penelitian ini. Tahapan dimulai dengan input data balita melalui User Interface (UI) yang berupa website yang menggunakan *HTML* dan *CSS Tailwind*. Data yang telah dimasukkan kemudian dikirimkan ke server yang menggunakan *Flask* untuk diproses lebih lanjut. Setelah data diproses, server mengirimkan data tersebut ke model machine learning untuk dilakukan klasifikasi status stunting balita. Hasil klasifikasi yang diperoleh kemudian dikembalikan dan ditampilkan di website. Selain itu, hasil klasifikasi bersama data input yang dimasukkan disimpan dalam database MySQL untuk pelacakan riwayat yang dapat diakses kembali oleh pengguna melalui fitur Riwayat pada website. Alur ini menegaskan bahwa sistem bekerja secara terstruktur, dimulai dari input data hingga penyajian hasil akhir dalam bentuk klasifikasi status stunting balita.

3.1.3 Desain Algoritma

Secara umum, penelitian ini merancang proses klasifikasi stunting melalui empat tahapan utama yang saling berkesinambungan. Keempat tahapan tersebut adalah normalisasi fitur, pembentukan kernel gabungan, pelatihan model *Support Vector Machine* (SVM), dan klasifikasi data uji. Seluruh tahapan ini bekerja secara berurutan layaknya sebuah alur produksi, di mana hasil dari satu tahap menjadi bahan baku bagi tahap berikutnya. Tujuan akhir dari alur ini adalah menghasilkan sebuah model prediktif yang mampu membedakan balita *stunting* dan normal secara konsisten. Data yang digunakan terbagi menjadi dua bagian, yaitu data latih untuk membangun *hyperplane* pemisah, dan data uji untuk menguji sejauh mana model dapat menebak dengan benar pada data yang belum pernah dilihatnya sebelumnya. Tabel 3.3 menyajikan tiga sampel balita yang mewakili kedua kelas sebagai titik awal ilustrasi komputasi.

Tabel 3. 3 Contoh Data Latih dan Data Uji

child_id	Jenis Kelamin	Umur Per Bulan	Tinggi (cm)	Berat (kg)	Pendidikan Ibu	Status Perawatan Makanan	Status Perawatan Kesehatan	Menerima vitamin A	Imunisasi Rotavirus 1
Balita 1	1	36	92	13.0	60	3	3	3	1
Balita 2	1	30	84	10.5	61	2	2	0	3
Balita 3	1	36	86	11	60	3	3	3	1

3.1.3.1 Normalisasi Data

Tahap normalisasi data menjadi fondasi penting sebelum memasuki proses pembentukan kernel. Penelitian ini menerapkan normalisasi *Z-score* untuk menyetarakan skala seluruh fitur numerik. Secara matematis, transformasi ini dirumuskan sebagai berikut:

$$Z = \frac{x - \mu}{\sigma} \quad (3.1)$$

Pada persamaan tersebut, Z adalah nilai hasil normalisasi, x adalah nilai asli dari fitur, μ (mu) adalah rata-rata fitur pada data latih, dan σ (sigma) adalah standar deviasi fitur pada data latih. Dengan transformasi ini, setiap fitur akan memiliki nilai rata-rata 0 dan standar deviasi 1, sehingga seluruh fitur berada dalam skala yang setara.

Langkah ini tidak dapat dilewati karena algoritma SVM, khususnya dengan kernel RBF, sangat bergantung pada perhitungan jarak Euclidean antar titik data. Apabila satu fitur memiliki rentang nilai yang jauh lebih lebar dibandingkan fitur lainnya, fitur dengan skala besar secara tidak proporsional akan mendominasi perhitungan jarak. Kondisi ini menyebabkan fitur-fitur lain kehilangan pengaruhnya dalam membentuk *hyperplane* pemisah, padahal secara klinis fitur tersebut mungkin memiliki hubungan yang kuat dengan status stunting. Normalisasi *Z-score* hadir untuk mengatasi ketimpangan skala tersebut, sehingga model tidak lagi memandang fitur berdasarkan besar kecilnya satuan pengukuran, melainkan berdasarkan pola hubungan antarfitur terhadap variabel target.

Penelitian ini menerapkan prosedur normalisasi dengan disiplin yang ketat. Nilai rata-rata (μ) dan standar deviasi (σ) untuk setiap fitur dihitung secara eksklusif dari data latih. Parameter statistik yang sama kemudian digunakan untuk mentransformasi data uji. Pemisahan yang tegas antara data latih dan data uji dalam proses normalisasi ini sangat krusial. Dalam implementasi di dunia nyata, model tidak pernah memiliki akses terhadap statistik data baru yang belum pernah

dilihatnya. Dengan menerapkan aturan ini, penelitian mensimulasikan skenario yang sesungguhnya, sehingga hasil evaluasi performa model menjadi lebih jujur dan dapat dipercaya.

Selain menyetarakan kontribusi fitur, normalisasi ini juga memberikan keuntungan dari sisi komputasi. Proses optimasi kuadratik pada SVM berjalan lebih stabil secara numerik ketika semua fitur berada dalam skala yang sebanding. Kondisi ini membantu algoritma mencapai titik optimum dengan lebih efisien dan menghindari potensi ketidakstabilan perhitungan yang dapat timbul akibat perbedaan skala yang ekstrem. Tabel 3.4 menampilkan hasil normalisasi untuk ketiga sampel balita. Terlihat bahwa seluruh fitur kini berada dalam skala yang sebanding, sehingga perhitungan jarak antar data menjadi lebih adil dan benar-benar mencerminkan tingkat kemiripan yang sesungguhnya antar balita.

Tabel 3. 4 Normalisasi Z-score

child_id	Jenis Kelamin	Umur Per Bulanan	Tinggi (cm)	Berat (kg)	Pendidikan Ibu	Status Perawatan Makanan	Status Perawatan Kesehatan	Menerima vitamin A	Imunisasi Rotavirus 1
Balita 1	0,0	1,0	1,0	1,0	0,0	1,0	1,0	1,0	0,0
Balita 2	0,0	0,0	0,0	0,0	1,0	0,0	0,0	0,0	1,0
Balita 3	0,0	1,0	0,333	0,083	0,0	0,0	0,0	1,0	0,0

3.1.3.2 Pembentukan Kernel Gabungan

Setelah data dinormalisasi, tahap selanjutnya adalah membangun fungsi kernel gabungan dengan mengacu pada kerangka *Multiple Kernel Learning* (MKL) yang telah dijelaskan secara teoritis pada Bab 2. Penelitian ini secara spesifik menggabungkan kernel linear dan kernel RBF menjadi satu fungsi kernel baru. Keputusan untuk menggabungkan kedua kernel ini didasarkan pada karakteristik

data stunting yang tidak seragam. Ada hubungan yang bersifat linear, misalnya korelasi langsung antara tinggi badan terhadap umur. Semakin tua usia anak dengan tinggi yang stagnan, semakin tinggi indikasi risiko stunting. Namun, ada pula interaksi yang bersifat non-linear dan kompleks, seperti efek sinergis antara rendahnya pendidikan ibu, ketidaklengkapan imunisasi, dan buruknya pola perawatan kesehatan. Kombinasi faktor-faktor ini sering kali memperkuat risiko stunting secara tidak proporsional, jauh melampaui efek penjumlahan linear masing-masing faktor. Di sinilah peran kernel RBF menjadi krusial.

Tiga parameter utama harus dioptimalkan secara bersamaan, yaitu bobot kernel (β), lebar fungsi RBF ($gamma$), dan koefisien regularisasi (C). Parameter β berfungsi sebagai pengatur keseimbangan. Jika β mendekati 1, model lebih mengandalkan hubungan linear. Jika β mendekati 0, model lebih mengandalkan hubungan non-linear. Parameter $gamma$ menentukan seberapa luas pengaruh sebuah titik data terhadap lingkungannya. Nilai $gamma$ yang kecil menghasilkan fungsi Gaussian yang lebar, yang berarti permukaan keputusan menjadi halus dan lebih general. Sebaliknya, nilai $gamma$ yang besar menghasilkan fungsi yang sempit, yang berarti model sangat peka terhadap detail lokal dan berisiko mengalami *overfitting*. Parameter C bertugas mengontrol toleransi model terhadap kesalahan klasifikasi pada data latih. Ketiga parameter ini saling terkait dan harus dieksplorasi secara simultan untuk menemukan konfigurasi terbaik.

Untuk memberikan gambaran yang lebih konkret, Tabel 3.5 menyajikan hasil perhitungan kernel untuk tiga sampel yang sudah dinormalisasi. Kernel linear dan RBF dihitung untuk setiap pasangan sampel, kemudian digabungkan dengan

bobot awal $\beta = 0.5$. Nilai kernel gabungan antara Balita 1 dan Balita 3 adalah 1,392, yang jauh lebih tinggi dibandingkan nilai antara Balita 2 dan Balita 3 yang hanya 0,184. Perbedaan yang mencolok ini mencerminkan bahwa Balita 1 dan Balita 3 memiliki profil antropometri dan pola pengasuhan yang sangat mirip, sementara Balita 2 memiliki profil yang berbeda. Secara geometris, nilai kernel yang tinggi berarti vektor-vektor fitur dari kedua balita tersebut memiliki arah yang sejalan (membentuk sudut kecil) dalam ruang fitur gabungan, sehingga kontribusinya terhadap keputusan klasifikasi akan lebih dominan.

Tabel 3. 5 Perhitungan Kernel Linear,RBF, dan Gabungan MKL

Pasangan Data	Kernel Linear (K_{linear})	Kernel RBF (K_{RBF})	Kernel Gabungan (K_{MKL}) ($\beta = 0.5$)
Balita 1 & Balita 2	0	0.135	0.068
Balita 1 & Balita 3	2.416	0.368	1.392
Balita 2 & Balita 3	0	0.368	0.184

3.1.3.3 Pelatihan Model SVM

Setelah matriks kernel gabungan terbentuk pada tahap sebelumnya, penelitian ini melanjutkan ke tahap pelatihan model SVM. Inti dari proses pelatihan ini adalah menyelesaikan permasalahan optimasi kuadratik cembung yang bertujuan menemukan hyperplane pemisah dengan margin terbesar, sekaligus meminimalkan kesalahan klasifikasi pada data latih. Karena penelitian ini menggunakan pendekatan kernel, optimasi dilakukan dalam bentuk dual dengan memanfaatkan pengali Lagrange (α). Pendekatan dual ini memiliki keunggulan tersendiri. Alih-alih mengoptimalkan vektor bobot secara langsung dalam ruang fitur berdimensi tinggi (yang bahkan mungkin berdimensi tak hingga pada kernel RBF), optimasi dual hanya melibatkan perhitungan hasil perkalian titik antar data

latih, yang dalam konteks ini digantikan oleh fungsi kernel. Konsep inilah yang dikenal sebagai kernel trick, yang memungkinkan model memanfaatkan kekuatan ruang fitur kompleks tanpa harus melakukan transformasi fitur secara eksplisit.

Secara sederhana, parameter regularisasi C pada persamaan optimasi berfungsi sebagai pengatur kompromi antara dua tujuan yang saling bertentangan. Di satu sisi, model ingin memaksimalkan lebar margin agar batas keputusan menjadi lebih general. Di sisi lain, model juga ingin meminimalkan kesalahan klasifikasi pada data latih. Nilai C yang besar memberikan penalti yang sangat tinggi terhadap setiap kesalahan klasifikasi, sehingga model akan berusaha keras untuk mengklasifikasikan seluruh data latih dengan benar. Akibatnya, hyperplane yang dihasilkan mungkin menjadi sangat kompleks dan sangat cocok dengan data latih (risiko overfitting). Sebaliknya, nilai C yang kecil membuat model lebih toleran terhadap kesalahan klasifikasi. Ia rela mengorbankan sedikit akurasi pada data latih demi mendapatkan margin yang lebih lebar, yang pada umumnya menghasilkan model dengan kemampuan generalisasi yang lebih baik terhadap data baru. Pemilihan nilai C yang tepat menjadi salah satu kunci keberhasilan model, dan itulah sebabnya penelitian ini mengoptimalkannya secara sistematis melalui grid search.

Keunggulan utama dari proses pelatihan SVM adalah sifat sparse dari solusi yang dihasilkan. Hanya data latih yang terletak tepat pada batas margin atau yang melanggar margin, yang disebut sebagai support vectors, yang memiliki nilai pengali Lagrange (α) tidak nol. Data-data lain yang berada jauh di dalam area kelasnya, dengan aman di sisi kanan hyperplane, memiliki nilai α sama dengan nol

dan sama sekali tidak memengaruhi model akhir. Dengan kata lain, dari ribuan data latih, hanya segelintir titik yang benar-benar menentukan posisi dan orientasi hyperplane. Sifat ini membawa konsekuensi yang sangat menguntungkan. Fungsi keputusan SVM hanya bergantung pada sekumpulan kecil titik data yang paling informatif. Akibatnya, proses prediksi untuk data baru menjadi sangat efisien secara komputasi, karena perhitungan hanya melibatkan kernel antara data baru dengan support vectors, bukan dengan seluruh data latih. Hal ini menjadikan SVM sangat cocok untuk diterapkan dalam sistem berbasis web yang memerlukan respons cepat.

Pada implementasi teknisnya, proses pelatihan ini dijalankan menggunakan pustaka Scikit-learn dengan opsi kernel precomputed. Opsi ini secara khusus dirancang untuk menerima matriks kernel kustom, seperti matriks gabungan MKL yang telah dibangun pada tahap sebelumnya. Dengan opsi ini, model tidak perlu menghitung kernel secara internal, melainkan langsung menggunakan matriks kemiripan yang telah disediakan. Hal ini memastikan bahwa pendekatan MKL dapat diintegrasikan secara mulus ke dalam kerangka kerja SVM standar tanpa modifikasi mendasar pada algoritma inti. Hasil dari proses pelatihan ini, yaitu kumpulan support vectors, nilai α , dan bias (b^*), kemudian menjadi fondasi bagi tahap klasifikasi data uji yang akan dijelaskan pada bagian selanjutnya.

3.1.3.4 Klasifikasi Data Uji

Tahap keempat menerapkan fungsi keputusan SVM untuk mengklasifikasikan data uji. Fungsi keputusan ini secara geometris menghitung jarak bertanda dari data uji terhadap hyperplane pemisah. Pada Tabel 3.6, Balita 3 (data uji) diklasifikasikan berdasarkan kontribusi dua support vector dari data latih,

yaitu Balita 1 yang berlabel normal dan Balita 2 yang berlabel stunting. Kontribusi masing-masing support vector dihitung sebagai hasil perkalian antara α , label kelas, dan nilai kernel gabungan antara data uji dan data latih. Balita 1 memberikan kontribusi positif sebesar 0,696 karena memiliki kemiripan kernel yang tinggi (1,392) dengan Balita 3, sementara Balita 2 memberikan kontribusi negatif sebesar -0,092 karena kemiripannya lebih rendah (0,184). Setelah ditambahkan dengan bias, total nilai fungsi keputusan ($f(x)$) yang dihasilkan adalah 0,704.

Karena nilai $f(x)$ tersebut positif (≥ 0), Balita 3 secara otomatis diklasifikasikan sebagai balita normal. Nilai absolut dari $f(x)$, yaitu 0,704, mencerminkan tingkat keyakinan model terhadap keputusannya—semakin besar nilai absolutnya, semakin jauh posisi data uji dari hyperplane, dan semakin tinggi keyakinan model. Ilustrasi ini menegaskan bahwa seluruh tahapan matematis dalam desain algoritma, mulai dari normalisasi hingga fungsi keputusan, telah terimplementasikan secara konsisten dan koheren sebelum dieksekusi ke dalam kode program.

Tabel 3. 6 Fungsi Keputusan Data Uji

Pasangan Data	Kernel Gabungan (K_{MKL})	α_i	y_i	$\alpha_i y_i K(x_i, x)$
Balita 1 & Balita 3	1.392	0.5	+1	$0.5 \times 1 \times 1.392 = 0.696$
Balita 2 & Balita 3	0,184	0.5	-1	$0.5 \times (-1) \times 0.184 = -0.092$
Total ($f(x)$)	$0.696 - 0.092 + 0.1 = 0.704$			

Keseluruhan alur komputasi ini diimplementasikan dalam bahasa Python dengan memanfaatkan pustaka *Scikit-learn* untuk pelatihan SVM dan pembentukan kernel, pustaka *Imbalanced-learn* untuk menyeimbangkan kelas melalui teknik SMOTENC, serta pustaka *Kneed* untuk analisis kepentingan fitur. Hasil dari setiap

tahapan inilah yang kemudian menjadi bahan evaluasi performa model pada Bab IV, di mana pengaruh setiap parameter dan fitur akan dianalisis secara kuantitatif.

3.1.4 Skenario Pengujian

Pengujian model SVM-MKL dilaksanakan dalam empat tahap yang saling melengkapi untuk memperoleh pemahaman komprehensif mengenai performa model dan kontribusi setiap fitur. Tahap pertama mencakup optimasi *hyperparameter* pada konfigurasi sembilan fitur yang dilanjutkan dengan evaluasi menyeluruh melalui validasi silang 10-*fold*. Tahap kedua berupa analisis kontribusi masing-masing fitur menggunakan *Permutation Feature Importance* dan seleksi fitur otomatis berbasis algoritma Kneedle. Tahap ketiga merupakan evaluasi model pada himpunan fitur yang telah terseleksi. Tahap keempat adalah pengujian pada himpunan fitur yang tidak terpilih sebagai konfirmasi empiris atas rendahnya kapasitas diskriminatif fitur-fitur tersebut. Keseluruhan konfigurasi diuji dengan prosedur yang identik untuk menjamin perbandingan kinerja yang adil.

3.1.4.1 K-Fold Cross Validation

K-Fold Cross Validation membagi keseluruhan data menjadi K himpunan bagian yang saling lepas dengan ukuran yang hampir sama. Pada setiap iterasi, satu himpunan digunakan sebagai data uji, sedangkan $K - 1$ himpunan lainnya digabungkan sebagai data latih. Prosedur ini diulang sebanyak K kali sehingga setiap himpunan pernah menjadi data uji tepat satu kali. Secara matematis, estimasi rata-rata performa model dinyatakan sebagai:

$$CV = \frac{1}{K} \sum_{i=1}^K Performance_i \quad (3.6)$$

Dengan K adalah jumlah *fold* dan $Performance_i$ adalah nilai metrik evaluasi pada iterasi ke- i . Nilai K yang digunakan dalam penelitian ini adalah 10. Pemilihan $K = 10$ didasarkan pada kemampuannya menjaga keseimbangan antara bias dan variansi estimasi performa. Nilai K yang terlalu kecil, misalnya $K = 2$, menghasilkan estimasi dengan bias rendah tetapi variansi tinggi karena ukuran data latih yang terbatas. Sebaliknya, K yang terlalu besar mendekati *leave-one-out cross validation* yang memiliki variansi rendah namun memerlukan biaya komputasi tinggi dan rentan terhadap *overfitting*. Dengan $K = 10$, setiap amatan dalam dataset memperoleh kesempatan menjadi data uji, sehingga estimasi akhir bersifat representatif terhadap populasi data (Tougui et al., 2021).

Penerapan validasi silang 10-*fold* juga memungkinkan pengukuran konsistensi model. Rata-rata setiap metrik menggambarkan performa model secara keseluruhan, sedangkan simpangan baku menunjukkan seberapa besar variasi performa ketika model dihadapkan pada subset data yang berbeda. Semakin kecil simpangan baku, semakin stabil model dalam memberikan prediksi. Stabilitas ini penting dalam konteks skrining kesehatan karena model akan digunakan pada populasi dengan karakteristik yang mungkin berbeda dari data pelatihan (Aprihartha & Idham, 2024).

3.1.4.2 Optimalisasi Hyperparameter

Sebelum evaluasi 10-fold dilaksanakan, setiap konfigurasi fitur menjalani proses pencarian hyperparameter terbaik menggunakan hold-out validation. Data

dipisahkan menjadi data latih dan data validasi dengan proporsi yang tetap sepanjang proses tuning. Pendekatan ini memungkinkan eksplorasi ruang parameter secara efisien tanpa melanggar prinsip dasar pemisahan data.

Grid search diterapkan untuk mengeksplorasi kombinasi nilai tiga hyperparameter utama MKL, yaitu C , γ , dan β . Setiap kombinasi diuji pada data validasi yang sama. Kriteria pemilihan parameter terbaik tidak bertumpu semata-mata pada F1-score tertinggi, melainkan juga mempertimbangkan keseimbangan antara presisi dan recall. Pertimbangan ini berangkat dari kesadaran bahwa dalam konteks skrining stunting, kesalahan positif palsu dan negatif palsu menimbulkan konsekuensi yang sama seriusnya. Model yang hanya mengoptimalkan F1-score dapat mengorbankan salah satu sisi, sehingga dipilih kombinasi yang memberikan perlindungan seimbang terhadap kedua jenis kesalahan tersebut.

Konfigurasi parameter terbaik yang diperoleh dari proses ini selanjutnya digunakan untuk melatih model final yang kemudian dievaluasi menggunakan validasi silang 10-fold. Prosedur yang serupa juga diterapkan pada subset fitur terpilih dan fitur tidak terpilih, sehingga setiap konfigurasi mendapatkan perlakuan yang setara.

3.1.4.3 Permutation Feature Importance

Setelah model sembilan fitur dilatih dan dievaluasi, kontribusi setiap fitur diukur menggunakan Permutation Feature Importance. Teknik ini bekerja dengan mengacak nilai satu fitur pada data uji, kemudian menghitung selisih antara F1-score awal dan F1-score setelah pengacakan. Pendekatan ini diperkenalkan oleh

Breiman (2001) dalam konteks Random Forest dan telah diadopsi secara luas untuk berbagai model machine learning karena bersifat model-agnostic, tidak bergantung pada asumsi distribusi data, dan mampu menangkap interaksi antarvariabel yang tidak terdeteksi oleh metode berbasis koefisien. Semakin besar penurunan F1-score setelah suatu fitur diacak, semakin penting fitur tersebut bagi model dalam menghasilkan prediksi yang akurat. Sebaliknya, fitur yang pengacakannya hampir tidak mengubah F1-score dianggap memberikan informasi yang tumpang tindih dengan fitur lain atau kurang memiliki kekuatan diskriminatif.

Penentuan fitur yang dipertahankan tidak dilakukan dengan menetapkan angka tetap secara manual, melainkan menggunakan pendekatan kalibrasi ambang otomatis yang adaptif terhadap data. Pendekatan ini bekerja dengan menganalisis selisih penurunan F1-score antarfitur yang telah diurutkan dari yang terbesar ke terkecil. Ketika selisih penurunan mulai mengecil secara drastis dan stabil pada nilai yang sangat rendah, hal ini menandakan bahwa fitur-fitur berikutnya sudah tidak memberikan kontribusi yang berarti dan dapat dianggap sebagai derau. Deteksi titik perubahan ini dilakukan secara otomatis menggunakan algoritma Kneedle, yang secara matematis menemukan titik dengan kelengkungan maksimum pada kurva penurunan (Satopaa et al., 2011). Setelah titik siku teridentifikasi, ambang batas ditetapkan pada celah alami antara fitur yang masih informatif dan fitur yang sudah menjadi derau. Pendekatan ini selaras dengan metode data-driven threshold yang direkomendasikan oleh Wang et al. (2023), di mana batas pemisah antara fitur relevan dan tidak relevan seharusnya diturunkan dari karakteristik data itu sendiri, bukan dari angka universal.

3.1.4.4 Evaluasi Subset Fitur

Himpunan fitur yang memenuhi kriteria tersebut selanjutnya digunakan untuk melatih ulang model. Sebelum evaluasi, dilakukan tuning *hyperparameter* menggunakan *grid search* dengan ruang pencarian yang identik dengan konfigurasi lengkap. Tuning ulang menjadi langkah penting karena reduksi dimensi fitur mengubah geometri ruang input dan distribusi data dalam ruang kernel. Permukaan fungsi objektif yang dihadapi model pun bergeser, sehingga parameter optimal untuk sembilan fitur belum tentu tetap optimal untuk jumlah fitur yang lebih sedikit. *Grid search* ulang memastikan bahwa model bekerja pada konfigurasi yang benar-benar sesuai dengan karakteristik data yang telah direduksi.

Setelah parameter terbaik diperoleh, model dievaluasi kembali dengan prosedur 10-*fold* yang sama. Performa kedua konfigurasi, sembilan fitur lengkap dan subset fitur terseleksi, kemudian dibandingkan untuk menilai dampak pengurangan fitur. Seluruh metrik evaluasi yang digunakan, yaitu *accuracy*, *precision*, *recall*, *F1-score*, dan *AUC*, seragam di seluruh tahap untuk menjaga konsistensi pengukuran.

3.1.4.5 Evaluasi Konfigurasi Fitur tidak Terpilih

Sebagai validasi tambahan, fitur-fitur yang tidak terpilih oleh algoritma Kneedle juga diuji sebagai konfigurasi mandiri. Fitur-fitur ini mencatat penurunan F1-score di bawah titik siku dan secara otomatis dikeluarkan dari model utama. Konfigurasi ini berfungsi sebagai kontrol negatif yang memperkuat validitas internal proses seleksi fitur. Apabila fitur-fitur yang dikeluarkan benar-benar tidak

memiliki kekuatan diskriminatif, maka model yang hanya mengandalkan fitur-fitur tersebut seharusnya menunjukkan performa yang jauh lebih rendah dibandingkan model yang menggunakan fitur-fitur hasil seleksi Kneedle.

Prosedur pengujian pada konfigurasi ini identik dengan konfigurasi lainnya, meliputi optimasi *hyperparameter* dengan *hold-out validation* dan evaluasi menggunakan validasi silang *10-fold*. Konsistensi prosedur memastikan bahwa setiap perbedaan performa yang muncul semata-mata disebabkan oleh perbedaan kandungan informasi dalam himpunan fitur, bukan oleh perbedaan metode evaluasi. Hasil dari ketiga konfigurasi ini selanjutnya dibandingkan secara langsung untuk menilai validitas proses seleksi fitur secara objektif. Temuan dari perbandingan ini akan saling mengonfirmasi dan memberikan gambaran yang utuh mengenai kontribusi sebenarnya dari setiap fitur terhadap klasifikasi stunting.

3.2 Implementasi

Tahap implementasi terdiri atas dua bagian utama, yaitu pembangunan antarmuka website dan penanaman algoritma SVM-MKL ke dalam sistem. Antarmuka website memungkinkan pengguna memasukkan data balita dan memperoleh hasil klasifikasi, sementara algoritma di sisi server memproses data masukan menggunakan model yang telah dilatih.

3.2.1 Implementasi Website

Implementasi website merupakan tahap penerapan rancangan sistem klasifikasi stunting berbasis web ke dalam aplikasi yang dapat dijalankan dan diuji langsung oleh pengguna. Sistem ini menggunakan HTML dan *Tailwind* CSS untuk

antarmuka, *Flask* API berbasis Python untuk *backend* yang memproses data menggunakan metode SVM-MKL, serta MySQL untuk menyimpan data dan hasil klasifikasi. Pengembangan dilakukan berdasarkan desain sistem yang telah dijelaskan pada bab sebelumnya guna menghasilkan aplikasi yang responsif, sederhana, dan efisien.

Gambar 3. 3 Halaman Klasifikasi

Pada Gambar 3.3 di atas, tampilan halaman ini dirancang untuk memfasilitasi proses input data balita yang diperlukan dalam klasifikasi status stunting. Halaman ini menyajikan form isian yang mencakup beberapa variabel penting, seperti umur, berat badan, tinggi badan, dan parameter fisik lainnya yang relevan dengan analisis. Setiap kolom input disediakan dengan *dropdown* menu atau field yang memungkinkan pengguna untuk memilih atau memasukkan data sesuai dengan kebutuhan. Setelah seluruh data terisi, pengguna dapat mengklik

tombol "Cek Klasifikasi" untuk memulai proses perhitungan klasifikasi menggunakan metode *Support Vector Machine* (SVM) dengan pendekatan *Multiple Kernel Learning* (MKL).

Selain itu, pada halaman ini terdapat fitur "Baca Cepat" yang memberikan penjelasan singkat mengenai cara penggunaan halaman tersebut dan petunjuk untuk melengkapi data jika ada kolom yang belum terisi. Bagian bawah halaman menampilkan status hasil klasifikasi, yang menunjukkan apakah data sudah diproses atau masih terdapat kekurangan dalam pengisian. Desain halaman ini dirancang untuk memastikan interaksi yang efisien antara pengguna dan sistem, dengan tujuan menghasilkan klasifikasi yang akurat dan relevan dalam mendeteksi status stunting pada balita secara cepat dan mudah.

RIWAYAT KLASIFIKASI				
Riwayat hasil Klasifikasi balita				
Kembali ke Form				
RINGKASAN				
TOTAL DATA				
13				
STATUS KONEKSI				
Berhasil dimuat				
NAVIGASI HALAMAN				
Jelajahi riwayat per halaman				
Menampilkan 1-5 dari 13 data				
POSISI SAAT INI				
Halaman 1 / 3				
TOTAL HALAMAN				
3 halaman				
Sebelumnya				
Berikutnya				
TABEL RIWAYAT				
Hasil klasifikasi tersimpan				
Klik tombol untuk memuat ulang data riwayat secara manual.				
Muat Ulang				
PROFIL ANAK	DATA FISIK	PERAWATAN	HASIL	WAKTU
JENIS KELAMIN 1	TINGGI DAN BERAT 88 cm Umur 24 bin Berat 10 kg	ASUPAN DAN KESEHATAN 2 Makanan 2 Kesehatan 2 Vitamin A 1 Rotavirus 1 Pendidikan Ibu 2	KLASIFIKASI STUNTING Status hasil tersimpan	CREATED AT 10 Jun 2026, 04.51 Riwayat klasifikasi
JENIS KELAMIN 1	TINGGI DAN BERAT 80 cm Umur 39 bin Berat 10 kg	ASUPAN DAN KESEHATAN 1 Makanan 1 Kesehatan 1 Vitamin A 0 Rotavirus 3 Pendidikan Ibu 15	KLASIFIKASI STUNTING Status hasil tersimpan	CREATED AT 08 Jun 2026, 19.50 Riwayat klasifikasi
JENIS KELAMIN 1	TINGGI DAN BERAT 150 cm Umur 40 bin Berat 84.99 kg	ASUPAN DAN KESEHATAN 3 Makanan 3 Kesehatan 3 Vitamin A 1 Rotavirus 1 Pendidikan Ibu 14	KLASIFIKASI NORMAL Status hasil tersimpan	CREATED AT 08 Jun 2026, 19.48 Riwayat klasifikasi
JENIS KELAMIN 1	TINGGI DAN BERAT 75 cm Umur 11 bin Berat 10 kg	ASUPAN DAN KESEHATAN 2 Makanan 2 Kesehatan 2 Vitamin A 2 Rotavirus 3 Pendidikan Ibu 3	KLASIFIKASI STUNTING Status hasil tersimpan	CREATED AT 26 Apr 2026, 23.37 Riwayat klasifikasi
JENIS KELAMIN 3	TINGGI DAN BERAT 104 cm Umur 42 bin Berat 17.7 kg	ASUPAN DAN KESEHATAN 3 Makanan 3 Kesehatan 3 Vitamin A 5 Rotavirus 3 Pendidikan Ibu 5	KLASIFIKASI NORMAL Status hasil tersimpan	CREATED AT 26 Apr 2026, 23.35 Riwayat klasifikasi

Gambar 3. 4 Halaman Riwayat Klasifikasi

Pada Gambar 3.4 yang menggambarkan halaman Riwayat Klasifikasi, halaman ini menampilkan tabel hasil klasifikasi status stunting balita yang telah dilakukan sebelumnya. Setiap baris pada tabel mencakup informasi seperti data balita, tanggal klasifikasi, dan status stunting (normal atau stunting). Pengguna dapat menavigasi antar halaman untuk melihat lebih banyak data klasifikasi yang telah diproses. Selain itu, terdapat ringkasan informasi di bagian kiri halaman yang menunjukkan jumlah data yang telah diproses dan jumlah halaman yang tersedia.

3.2.2 Implementasi Algoritma

Peneliti mengimplementasikan algoritma SVM-MKL dalam bahasa Python dengan memanfaatkan pustaka *Scikit-learn*, *imbalanced-learn*, dan *Kneed*. Fungsi `linear_kernel` dan `mkl_kernel` membentuk kernel gabungan yang menggabungkan kernel linear dan RBF dengan bobot beta. Proses tuning dijalankan oleh `tune_mkl_holdout` yang mengeksplorasi 27 kombinasi C, gamma, dan beta menggunakan *hold-out validation*. Setiap kombinasi dinormalisasi, dilatih, dan dievaluasi pada data validasi yang sama. Hasil dari seluruh kombinasi kemudian diseleksi oleh `select_best_combination` yang memilih konfigurasi dengan *F1-score* tertinggi. Jika terdapat beberapa kandidat dengan *F1-score* yang hampir setara, fungsi ini memilih kombinasi yang memiliki selisih paling kecil antara presisi dan *recall*, sehingga keseimbangan kedua metrik tetap terjaga.

Evaluasi menyeluruh dilakukan oleh `evaluate_mkl_cv` yang menjalankan validasi silang 10-fold. Pada setiap lipatan, data latih diseimbangkan dengan SMOTENC, dinormalisasi, dan matriks kernel MKL dihitung. Model SVM dilatih dengan kernel precomputed, lalu metrik akurasi, presisi, recall, *F1-score*, dan AUC

dicatat. Rata-rata dan simpangan baku dari kesepuluh lipatan dikembalikan sebagai estimasi akhir performa model. Tingkat kepentingan fitur diukur oleh *permutation_importance*. Fungsi ini mencatat baseline *F1-score*, kemudian mengacak setiap kolom fitur secara bergantian sebanyak lima kali. Rata-rata penurunan *F1-score* akibat pengacakan dicatat sebagai skor kepentingan. Fitur-fitur dengan penurunan besar dianggap penting, sedangkan yang kecil dianggap kurang berkontribusi.

Algoritma Kneedle mendeteksi titik siku pada kurva penurunan dan secara otomatis memilih fitur-fitur yang dipertahankan. Pada alur utama, model dilatih pada sembilan fitur, dievaluasi, dan disimpan. Setelah fitur-fitur terseleksi diperoleh, tuning ulang dilakukan pada subset fitur, diikuti dengan evaluasi dan penyimpanan model subset. Kode inti dari implementasi ini disajikan sebagai satu kesatuan yang mencakup seluruh fungsi dan alur pemanggilannya.

```
import pandas as pd
import numpy as np
from sklearn.preprocessing import StandardScaler
from sklearn.model_selection import KFold, train_test_split
from sklearn.metrics import accuracy_score, precision_score,
recall_score, f1_score, roc_auc_score
from sklearn.svm import SVC
from sklearn.metrics.pairwise import rbf_kernel
from imblearn.over_sampling import SMOTENC
import joblib
from kneed import KneeLocator

# Fungsi kernel
def linear_kernel(X1, X2):
    return np.dot(X1, X2.T)

def mkl_kernel(X1, X2, beta, gamma):
    K_lin = linear_kernel(X1, X2)
    K_rbf = rbf_kernel(X1, X2, gamma=gamma)
    return beta * K_lin + (1 - beta) * K_rbf

# Tuning dengan hold-out validation
def tune_mkl_holdout(X_raw, y, cat_cols_list, fitur_names,
                    C_vals, gamma_vals, beta_vals,
                    use_smotenc=True, test_size=0.2):
```

```

        X_train_raw, X_val_raw, y_train, y_val = train_test_split(
            X_raw, y, test_size=test_size, random_state=42,
            stratify=y
        )
        cat_indices = [i for i, col in enumerate(fitur_names) if col
            in cat_cols_list]
        all_results = []
        for C in C_vals:
            for gamma in gamma_vals:
                for beta in beta_vals:
                    if use_smotenc:
                        smote =
SMOTENC(categorical_features=cat_indices, random_state=42)
                        X_tr_res, y_tr_res =
smote.fit_resample(X_train_raw, y_train)
                    else:
                        X_tr_res, y_tr_res = X_train_raw, y_train
                        scaler = StandardScaler()
                        X_tr = scaler.fit_transform(X_tr_res)
                        X_val = scaler.transform(X_val_raw)
                        K_train = mkl_kernel(X_tr, X_tr, beta, gamma)
                        K_val = mkl_kernel(X_val, X_tr, beta, gamma)
                        svm = SVC(kernel='precomputed', C=C)
                        svm.fit(K_train, y_tr_res)
                        y_pred = svm.predict(K_val)
                        acc = accuracy_score(y_val, y_pred)
                        prec = precision_score(y_val, y_pred)
                        rec = recall_score(y_val, y_pred)
                        f1 = f1_score(y_val, y_pred)
                        all_results.append({
                            'C': C, 'gamma': gamma, 'beta': beta,
                            'accuracy': acc, 'precision': prec,
                            'recall': rec, 'f1': f1
                        })
        return all_results

# Pemilihan kombinasi terbaik berdasarkan F1-score dan
keseimbangan presisi-recall
def select_best_combination(results, f1_tolerance=0.001):
    sorted_results = sorted(results, key=lambda x: x['f1'],
        reverse=True)
    best_f1 = sorted_results[0]['f1']
    candidates = [r for r in sorted_results if abs(r['f1'] -
        best_f1) <= f1_tolerance]
    if len(candidates) == 1:
        return candidates[0]
    return min(candidates, key=lambda x: abs(x['precision'] -
        x['recall']))

# Evaluasi 10-fold cross validation
def evaluate_mkl_cv(X_raw, y, cat_cols_list, fitur_names, C,
    gamma, beta, use_smotenc=True):
    kf = KFold(n_splits=10, shuffle=True, random_state=42)
    cat_indices = [i for i, col in enumerate(fitur_names) if col
        in cat_cols_list]
    metrics = []

```

```

    for train_idx, test_idx in kf.split(X_raw):
        X_train_raw, X_test_raw = X_raw.iloc[train_idx],
X_raw.iloc[test_idx]
        y_train, y_test = y.iloc[train_idx], y.iloc[test_idx]
        if use_smotenc:
            smote = SMOTENC(categorical_features=cat_indices,
random_state=42)
            X_train_res, y_train_res =
smote.fit_resample(X_train_raw, y_train)
        else:
            X_train_res, y_train_res = X_train_raw, y_train
            scaler = StandardScaler()
            X_train = scaler.fit_transform(X_train_res)
            X_test = scaler.transform(X_test_raw)
            K_train = mkl_kernel(X_train, X_train, beta, gamma)
            K_test = mkl_kernel(X_test, X_train, beta, gamma)
            svm = SVC(kernel='precomputed', C=C)
            svm.fit(K_train, y_train_res)
            y_pred = svm.predict(K_test)
            dec = svm.decision_function(K_test)
            metrics.append({
                'acc': accuracy_score(y_test, y_pred),
                'prec': precision_score(y_test, y_pred),
                'rec': recall_score(y_test, y_pred),
                'f1': f1_score(y_test, y_pred),
                'auc': roc_auc_score(y_test, dec)
            })
    df_metrics = pd.DataFrame(metrics)
    return df_metrics.mean(), df_metrics.std()

# Permutation feature importance
def permutation_importance(model, X_scaled, y_true, kernel_func,
beta, gamma, n_repeats=5):
    K_full = kernel_func(X_scaled, X_scaled, beta, gamma)
    baseline_f1 = f1_score(y_true, model.predict(K_full))
    importance = {}
    for col in range(X_scaled.shape[1]):
        drops = []
        for _ in range(n_repeats):
            X_perm = X_scaled.copy()
            np.random.shuffle(X_perm[:, col])
            K_perm = kernel_func(X_perm, X_scaled, beta, gamma)
            drops.append(baseline_f1 - f1_score(y_true,
model.predict(K_perm)))
        importance[col] = np.mean(drops)
    return importance, baseline_f1

# Alur utama
X_9 = df[fitur_9]
all_results_9 = tune_mkl_holdout(X_9, y, cat_cols, fitur_9,
C_vals, gamma_vals, beta_vals, use_smotenc)
best_comb_9 = select_best_combination(all_results_9)
C_best, gamma_best, beta_best = best_comb_9['C'],
best_comb_9['gamma'], best_comb_9['beta']

```

```

mean_9, std_9 = evaluate_mkl_cv(X_9, y, cat_cols, fitur_9,
C_best, gamma_best, beta_best, use_smotenc)

# Simpan model 9 fitur
scaler_9 = StandardScaler()
X_9_scaled = scaler_9.fit_transform(X_9_resampled)
K_9 = mkl_kernel(X_9_scaled, X_9_scaled, beta_best, gamma_best)
svm_9 = SVC(kernel='precomputed', C=C_best).fit(K_9,
y_9_resampled)
joblib.dump(scaler_9, 'scaler.pkl')
joblib.dump(svm_9, 'svm_mkl_model.pkl')

# Feature importance dan seleksi fitur
imp, baseline = permutation_importance(svm_9, X_9_scaled,
y_9_resampled, mkl_kernel, beta_best, gamma_best)
imp_df = pd.DataFrame({'Fitur': fitur_9, 'Penurunan_F1': [imp[i]
for i in range(len(fitur_9))]}).sort_values('Penurunan_F1',
ascending=False)
penurunan_vals = imp_df['Penurunan_F1'].values
kl = KneeLocator(range(1, len(penurunan_vals)+1),
penurunan_vals, curve='convex', direction='decreasing')
elbow_idx = kl.knee
if elbow_idx is not None:
    selected_features = imp_df.head(elbow_idx)['Fitur'].tolist()
else:
    selisih = np.diff(penurunan_vals)
    elbow_idx_manual = np.argmax(selisih) + 1
    selected_features = imp_df.head(elbow_idx_manual +
1)['Fitur'].tolist()

# Tuning ulang dan evaluasi subset fitur
X_sub = df[selected_features]
all_results_sub = tune_mkl_holdout(X_sub, y, cat_cols,
selected_features, C_vals, gamma_vals, beta_vals, use_smotenc)
best_comb_sub = select_best_combination(all_results_sub)
C_sub, gamma_sub, beta_sub = best_comb_sub['C'],
best_comb_sub['gamma'], best_comb_sub['beta']
mean_sub, std_sub = evaluate_mkl_cv(X_sub, y, cat_cols,
selected_features, C_sub, gamma_sub, beta_sub, use_smotenc)

# Simpan model subset
scaler_sub = StandardScaler()
X_sub_scaled = scaler_sub.fit_transform(X_sub_resampled)
K_sub = mkl_kernel(X_sub_scaled, X_sub_scaled, beta_sub,
gamma_sub)
svm_sub = SVC(kernel='precomputed', C=C_sub).fit(K_sub,
y_sub_resampled)
joblib.dump(scaler_sub, 'scaler_subset.pkl')
joblib.dump(svm_sub, 'svm_mkl_subset.pkl')

```

BAB IV

HASIL DAN PEMBAHASAN

4.1 Hasil Pengujian

Pengujian dilaksanakan secara bertahap sesuai dengan desain yang telah ditetapkan pada Bab III. Seluruh konfigurasi fitur dievaluasi dengan prosedur yang seragam agar perbandingan kinerja berlangsung secara adil.

4.1.1 Hasil Optimalisasi Hyperparameter

Pencarian hyperparameter terbaik untuk konfigurasi sembilan fitur dijalankan menggunakan *hold-out validation*. Data dipisahkan menjadi data latih dan data validasi dengan proporsi tetap sepanjang proses *tuning*. *Grid search* mengeksplorasi 27 kombinasi dari tiga parameter, yaitu C , γ , dan β . Ruang pencarian yang digunakan disajikan pada Tabel 4.1.

Tabel 4. 1 Ruang Pencarian Hyperparameter

Parameter	Nilai yang Diuji
C	0.1, 1, 10
γ	0.01, 0.1, 1
β	0.5, 0.6, 0.7

Rentang nilai pada Tabel 4.1 dipilih dengan mempertimbangkan rekomendasi Hsu et al. (2003) yang menyarankan penggunaan rentang eksponensial karena pengaruh C dan γ bersifat multiplikatif. Rentang 0.1 hingga 10 untuk C mencakup tiga tingkat regularisasi yang berbeda. Rentang 0.01 hingga 1 untuk γ mencakup fungsi Gaussian yang lebar, sedang, dan sempit. Untuk β , nilai 0.5, 0.6, dan 0.7 dipilih untuk menguji kontribusi kernel linear dari

bobot yang seimbang hingga sedikit lebih dominan. Dalam kerangka MKL, bobot optimal sering kali berada di sekitar nilai tengah ketika data mengandung pola linear dan non linear yang sama kuat (Gönen & Alpaydin, 2011).

Dari ruang pencarian tersebut diperoleh 27 kombinasi. Setiap kombinasi dievaluasi menggunakan rata rata F1 score dari kesepuluh lipatan. Hasil lengkap disajikan pada table 4.2.

Tabel 4. 2 Hasil Grid Search 27 Kombinasi Hyperparameter

Kombinasi	C	gamma	beta	F1 score	precision	recall	akurasi
1	0.1	0.01	0.5	0.8557	0.9105	0.8345	0.8708
2	0.1	0.01	0.6	0.8533	0.9063	0.8345	0.8689
3	0.1	0.01	0.7	0.8524	0.9049	0.8345	0.8683
4	0.1	0.1	0.5	0.9275	0.9613	0.9123	0.9361
5	0.1	0.1	0.6	0.9225	0.9581	0.9066	0.9317
6	0.1	0.1	0.7	0.9134	0.9561	0.8925	0.9232
7	0.1	1	0.5	0.9011	0.9387	0.8883	0.9128
8	0.1	1	0.6	0.8945	0.9302	0.8854	0.9072
9	0.1	1	0.7	0.8887	0.9256	0.8798	0.9021
10	1	0.01	0.5	0.9151	0.9548	0.8967	0.9249
11	1	0.01	0.6	0.9060	0.9486	0.8868	0.9167
12	1	0.01	0.7	0.9002	0.9453	0.8798	0.9114
13	1	0.1	0.5	0.9670	0.9799	0.9632	0.9715
14	1	0.1	0.6	0.9637	0.9784	0.9590	0.9686
15	1	0.1	0.7	0.9596	0.9727	0.9576	0.9651
16	1	1	0.5	0.9497	0.9654	0.9477	0.9565
17	1	1	0.6	0.9423	0.9570	0.9434	0.9501
18	1	1	0.7	0.9382	0.9553	0.9378	0.9465
19	10	0.01	0.5	0.9711	0.9855	0.9646	0.9750
20	10	0.01	0.6	0.9637	0.9784	0.9590	0.9686
21	10	0.01	0.7	0.9629	0.9783	0.9576	0.9678
22	10	0.1	0.5	0.9802	0.9886	0.9774	0.9829
23	10	0.1	0.6	0.9794	0.9871	0.9774	0.9822
24	10	0.1	0.7	0.9777	0.9857	0.9760	0.9808
25	10	1	0.5	0.9637	0.9716	0.9661	0.9688
26	10	1	0.6	0.9621	0.9715	0.9632	0.9673
27	10	1	0.7	0.9563	0.9658	0.9590	0.9624

Pemilihan parameter terbaik menggunakan *F1-score* sebagai acuan utama, bukan akurasi. Keputusan ini didasarkan pada karakteristik dataset yang digunakan. Dataset IFLS-5 memiliki distribusi kelas yang tidak seimbang, dengan jumlah

sampel stunting yang lebih banyak dibandingkan sampel normal. Pada kondisi ketidakseimbangan kelas, akurasi dapat memberikan gambaran yang menyesatkan. Model yang hanya menebak kelas mayoritas secara terus-menerus dapat memperoleh akurasi yang tampak tinggi, padahal sama sekali tidak mampu mendeteksi kelas minoritas. *F1-score* dipilih karena metrik ini merupakan rata-rata harmonis dari presisi dan *recall*. Rata-rata harmonis memberikan penalti yang lebih besar ketika salah satu komponennya bernilai rendah. Dengan demikian, *F1-score* hanya akan tinggi jika presisi dan *recall* sama-sama tinggi. Karakteristik ini menjadikan *F1-score* sebagai metrik yang lebih jujur dalam menilai performa model pada data dengan kelas yang tidak seimbang (Chicco & Jurman, 2020).

Dalam konteks skrining stunting, keseimbangan antara presisi dan *recall* sangat penting. Presisi yang tinggi berarti model jarang mengklasifikasikan anak normal sebagai stunting, sehingga orang tua tidak menerima peringatan yang keliru dan fasilitas kesehatan tidak dibebani rujukan yang tidak perlu. *Recall* yang tinggi berarti model mampu menemukan sebagian besar kasus stunting yang sesungguhnya, sehingga hanya sedikit anak yang luput dari deteksi. Kedua jenis kesalahan ini membawa konsekuensi serius. Positif palsu dapat mengikis kepercayaan masyarakat terhadap program skrining, sementara negatif palsu dapat menghilangkan kesempatan anak untuk mendapatkan intervensi gizi pada masa pertumbuhan kritis. *F1-score* yang menyeimbangkan keduanya menjadi pilihan yang tepat karena mencerminkan performa model secara lebih adil dibandingkan metrik tunggal lainnya (Hicks et al., 2022).

Penerapan *F1-score* sebagai metrik utama dalam pemilihan parameter juga terlihat pada proses pengambilan keputusan di Tabel 4.2. Kombinasi ke-22 dengan $C = 10$, $\gamma = 0.1$, dan $\beta = 0.5$ mencatat *F1-score* tertinggi sebesar 0.9829. Namun, kombinasi ke-23 dengan $C = 10$, $\gamma = 0.1$, dan $\beta = 0.6$ dipilih sebagai konfigurasi optimal karena menawarkan keseimbangan yang lebih baik antara presisi dan *recall*. Kombinasi ini mencatat presisi 0.9871 dan *recall* 0.9774. Presisi yang lebih tinggi berarti model lebih jarang mengklasifikasikan anak normal sebagai stunting. Meskipun *F1-score* kombinasi ke-23 sedikit lebih rendah, selisihnya hanya 0.0007, sebuah pengorbanan yang sangat kecil untuk memperoleh perlindungan lebih baik terhadap kesalahan positif palsu. Keputusan ini menegaskan bahwa *F1-score* digunakan sebagai acuan utama, namun tetap diinterpretasikan bersama dengan presisi dan *recall* secara terpisah untuk memastikan keseimbangan yang optimal antara kedua metrik tersebut.

Tabel 4.2 juga memperlihatkan pengaruh masing masing parameter terhadap performa model. Peningkatan nilai C secara konsisten mendorong kenaikan akurasi dan *F1-score*. Pada kelompok $C = 10$, seluruh kombinasi mencatat akurasi di atas 0.95. Parameter γ menunjukkan pengaruh yang tidak linier. $\gamma = 0.1$ memberikan hasil terbaik di semua kelompok C , sementara $\gamma = 1$ cenderung menurunkan presisi. Parameter β menunjukkan pengaruh yang lebih halus. Pada kelompok $C = 10$ dan $\gamma = 0.1$, $\beta = 0.5$ menghasilkan *recall* yang sedikit lebih tinggi, sedangkan $\beta = 0.6$ memberikan presisi yang lebih baik. Inilah yang menjadi dasar pemilihan konfigurasi ke 23 sebagai parameter terbaik.

4.1.2 Hasil Pengujian K-Fold

Model dengan parameter terbaik $C = 10$, $\gamma = 0.1$, dan $\beta = 0.6$ dievaluasi menggunakan validasi silang 10-fold. Prosedur ini membagi data menjadi sepuluh himpunan yang saling lepas. Pada setiap putaran, sembilan himpunan digunakan sebagai data latih dan satu himpunan sebagai data uji. Rata-rata dan simpangan baku dari kesepuluh lipatan menjadi estimasi akhir performa model. Rincian metrik disajikan pada Tabel 4.3.

Tabel 4. 3 Hasil Pengujian K-fold

Fold	Accuracy	Precision	Recall	F1 Score	AUC
1	0.9901	0.9971	0.9858	0.9914	0.9979
2	0.9752	0.9750	0.9832	0.9791	0.9953
3	0.9818	0.9833	0.9860	0.9847	0.9974
4	0.9868	0.9914	0.9857	0.9885	0.9986
5	0.9752	0.9832	0.9751	0.9791	0.9978
6	0.9670	0.9760	0.9644	0.9701	0.9975
7	0.9752	0.9854	0.9712	0.9782	0.9979
8	0.9769	0.9796	0.9796	0.9796	0.9971
9	0.9752	0.9860	0.9724	0.9791	0.9970
10	0.9785	0.9810	0.9837	0.9824	0.9986
Rata-rata $\pm$ SD	0.978 ± 0.007	0.984 ± 0.007	0.979 ± 0.008	0.981 ± 0.006	0.998 ± 0.001

Model mencapai rata rata accuracy 0.978. Artinya, 97,8% dari seluruh sampel diklasifikasikan dengan benar. Tingkat kesalahan hanya sebesar 2,2%. Simpangan baku accuracy yang hanya 0.007 menegaskan bahwa performa model seragam di seluruh fold. Tidak ada satu fold pun yang mengalami penurunan akurasi secara drastis.

Precision rata rata tercatat 0.984 dari seluruh balita yang diprediksi stunting oleh model, 98,4% di antaranya benar benar mengalami stunting. False positive, yaitu balita normal yang salah diklasifikasikan sebagai stunting, hanya berjumlah sekitar 1,6% dari keseluruhan prediksi stunting. Angka ini penting dalam praktik

skrining. Setiap kesalahan positif palsu dapat menimbulkan kecemasan yang tidak perlu pada orang tua dan membebani fasilitas kesehatan dengan rujukan yang tidak dibutuhkan.

Recall rata rata sebesar 0.979 menunjukkan bahwa model berhasil mendeteksi 97,9% dari seluruh kasus stunting yang sesungguhnya ada. False negative, yaitu balita stunting yang tidak terdeteksi, hanya sekitar 2,1%. Setiap kesalahan negatif palsu mewakili seorang anak yang kehilangan kesempatan untuk mendapatkan intervensi gizi dini. Dampaknya dapat bersifat permanen terhadap pertumbuhan fisik dan kognitif.

Keseimbangan antara precision dan recall tercermin dalam F1 score rata rata 0.981. F1 score merupakan rata rata harmonis dari kedua metrik tersebut. Nilai ini hanya akan tinggi jika keduanya sama sama tinggi. Model tidak mengorbankan salah satu sisi secara berlebihan. Karakteristik ini sangat diharapkan dalam sistem deteksi dini, di mana false positive maupun false negative sama sama harus diminimalkan.

AUC rata rata 0.998 dengan simpangan baku 0.001 memperlihatkan kemampuan diskriminasi yang hampir sempurna. AUC mengukur seberapa baik model membedakan kelas stunting dan normal tanpa bergantung pada ambang batas tertentu. Model mampu menempatkan balita stunting pada peringkat probabilitas yang lebih tinggi daripada balita normal dalam 99,8% kasus. Bahkan pada fold ke 6 yang memiliki accuracy dan recall terendah, AUC tetap bertahan di 0.998. Penurunan accuracy pada fold tersebut lebih disebabkan oleh pergeseran titik

potong klasifikasi, bukan oleh melemahnya kemampuan model dalam mengurutkan risiko.

4.1.3 Hasil *Feature Importance*

Tingkat kepentingan setiap fitur dalam model sembilan fitur diukur menggunakan *Permutation Feature Importance*. *Baseline F1-score* model sebelum pengacakan adalah 0.9891. Hasil pengukuran disajikan pada Tabel 4.4.

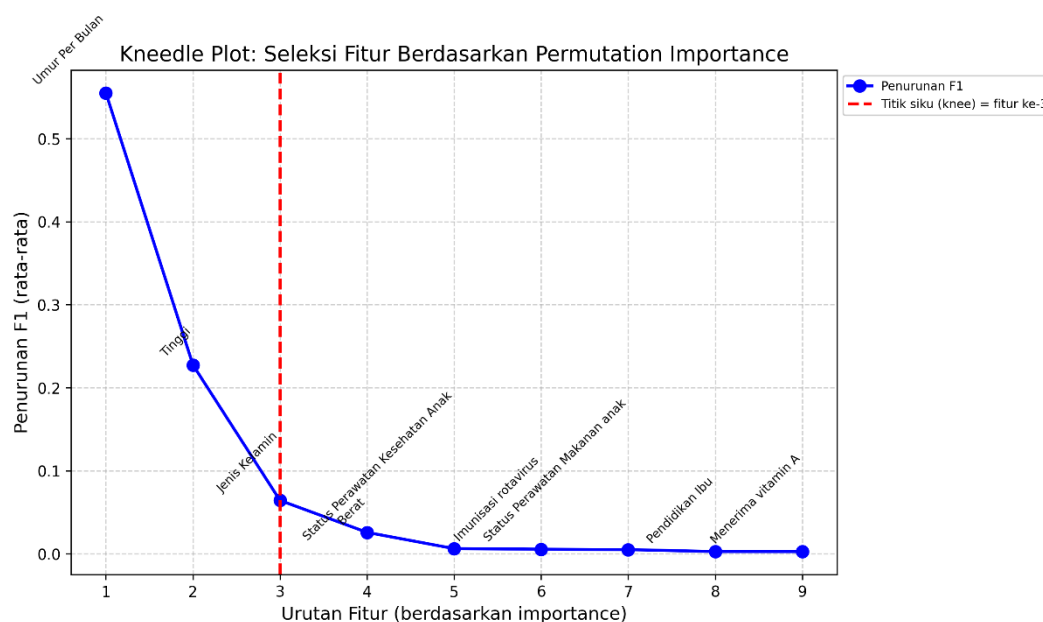
Tabel 4. 4 Hasil *Permutation Feature Importance*

Fitur	Penurunan F1
Umur Per Bulan	0.5545
Tinggi	0.2250
Jenis Kelamin	0.0643
Berat	0.0260
Status Perawatan Kesehatan Anak	0.0068
Status Perawatan Makanan anak	0.0056
Imunisasi rotavirus	0.0047
Pendidikan Ibu	0.0032
Menerima vitamin A	0.0030

Tabel 4.4 memperlihatkan kesenjangan yang sangat tajam antara tiga fitur teratas dan sisanya. Fitur Umur Per Bulan sendirian bertanggung jawab atas penurunan *F1-score* sebesar 0.5560. Angka ini berarti bahwa kemampuan model mendeteksi stunting menurun drastis ketika informasi umur dihilangkan. Lebih dari setengah kemampuannya lenyap. Temuan ini konsisten dengan definisi klinis stunting yang didasarkan pada *Height-for-Age Z-score* (HAZ). Umur menjadi sumbu utama yang menentukan posisi seorang anak dalam kurva pertumbuhan. Tinggi badan tidak dapat dinilai secara klinis tanpa informasi umur.

Fitur Tinggi mencatat penurunan 0.2246 dan menempati urutan kedua. Bersama Umur Per Bulan dan Jenis Kelamin, ketiga fitur ini adalah komponen utama HAZ. Algoritma Kneedle mendeteksi titik siku tepat pada fitur Jenis

Kelamin dengan penurunan $F1$ -score sebesar 0.0655. Keputusan ini menempatkan tiga fitur pertama sebagai himpunan terseleksi. Enam fitur berikutnya dikeluarkan. Deteksi titik siku terjadi pada fitur ketiga karena gradien penurunan antara Tinggi dan Jenis Kelamin masih cukup besar. Sebaliknya, gradien penurunan antara Jenis Kelamin dan Berat sudah melandai secara signifikan. Algoritma Kneedle secara matematis mendeteksi perubahan kelengkungan ini sebagai batas alami antara fitur yang informatif dan fitur yang hanya menyumbang derau.



Gambar 4. 1 Grafik Algoritma Kneedle Seleksi Fitur

Gambar 4.1 menyajikan kurva penurunan $F1$ -score yang telah diurutkan dari nilai tertinggi ke terendah. Sumbu horizontal menunjukkan urutan fitur berdasarkan peringkat kepentingan, sedangkan sumbu vertikal menunjukkan besarnya penurunan $F1$ -score setelah fitur diacak. Titik siku yang terdeteksi oleh algoritma Kneedle ditandai dengan garis vertikal putus-putus tepat pada fitur Jenis Kelamin. Fitur-fitur di sebelah kiri titik siku, yaitu Umur Per Bulan, Tinggi, dan

Jenis Kelamin, dipertahankan sebagai himpunan terseleksi. Sementara itu, fitur-fitur di sebelah kanan titik siku dikeluarkan karena penurunannya sudah melandai dan tidak lagi memberikan kontribusi yang berarti. Pola visual ini menegaskan bahwa proses seleksi fitur tidak didasarkan pada asumsi subjektif, melainkan pada karakteristik intrinsik data yang terukur secara kuantitatif.

4.1.4 Hasil Pengujian Subset Fitur

Tiga fitur yang dipertahankan oleh algoritma Kneedle, yaitu Umur Per Bulan, Tinggi, dan Jenis Kelamin, selanjutnya digunakan untuk melatih ulang model. Reduksi dari sembilan dimensi menjadi tiga dimensi bukan sekadar penyederhanaan jumlah variabel. Perubahan ini mengubah struktur ruang input secara fundamental. Data yang semula tersebar dalam ruang berdimensi tinggi kini menempati ruang yang jauh lebih ringkas. Jarak antartitik data yang sebelumnya renggang menjadi lebih padat, dan distribusi data dalam ruang kernel pun bergeser. Konsekuensinya, permukaan fungsi objektif yang dihadapi oleh algoritma optimasi tidak lagi sama dengan sebelumnya. Parameter C , γ , dan β yang optimal untuk konfigurasi sembilan fitur belum tentu tetap optimal untuk konfigurasi tiga fitur. Grid search ulang menjadi langkah yang tidak dapat dilewati untuk memastikan model bekerja pada konfigurasi yang benar-benar sesuai dengan karakteristik data yang telah direduksi.

Tuning ulang dijalankan menggunakan prosedur hold-out validation yang identik dengan konfigurasi awal. Parameter terbaik yang diperoleh adalah $C = 10$, $\gamma = 1.0$, dan $\beta = 0.5$. Ada satu perubahan yang mencolok: nilai γ bergeser dari 0.1 menjadi 1.0. Pergeseran ini memiliki makna yang penting. Ketika

ruang fitur masih berdimensi sembilan, data tersebar secara renggang sehingga model memerlukan fungsi Gaussian yang lebar agar satu titik data dapat menjangkau area yang cukup luas. Sebaliknya, ketika ruang fitur mengecil menjadi tiga dimensi, data menjadi lebih rapat dan setiap titik sudah saling berdekatan secara alami. Dalam kondisi ini, fungsi Gaussian yang lebih sempit justru lebih efektif karena model dapat menangkap perbedaan halus antarkasus tanpa harus menyebarkan pengaruh terlalu lebar. Ini adalah respons adaptif model terhadap penyederhanaan ruang fitur, yang menunjukkan bahwa model tidak kaku terhadap perubahan dimensi, melainkan mampu menyesuaikan kompleksitasnya sesuai dengan data yang dihadapi.

Evaluasi menggunakan validasi silang 10-fold kemudian dilakukan dengan parameter terbaik tersebut. Rincian metrik pada setiap fold disajikan pada Tabel 4.5.

Tabel 4. 5 Hasil Pengujian K-fold Subset Fitur

Fold	Accuracy	Precision	Recall	F1-Score	AUC
1	0.9934	0.9943	0.9943	0.9943	0.9999
2	0.9884	0.9972	0.9832	0.9901	0.9999
3	0.9950	1.0000	0.9916	0.9958	1.0000
4	0.9934	1.0000	0.9885	0.9942	1.0000
5	0.9901	1.0000	0.9834	0.9916	0.9999
6	0.9950	0.9970	0.9941	0.9955	1.0000
7	0.9917	0.9942	0.9914	0.9928	0.9998
8	0.9901	0.9971	0.9854	0.9912	0.9999
9	0.9950	0.9972	0.9945	0.9959	1.0000
10	0.9967	0.9973	0.9973	0.9973	1.0000
Rata-rata ± SD	0.993 ± 0.003	0.997 ± 0.002	0.990 ± 0.005	0.994 ± 0.002	1.000 ± 0.000

Model tiga fitur menunjukkan peningkatan yang sangat meyakinkan di seluruh metrik. Akurasi naik dari 0.978 menjadi 0.993, presisi naik dari 0.984 menjadi 0.997, *recall* naik dari 0.979 menjadi 0.990, *F1-score* naik dari 0.981 menjadi 0.994, dan *AUC* mencapai nilai sempurna 1.000. Peningkatan yang terjadi

cukup signifikan, terutama pada presisi yang kini hampir menyentuh kesempurnaan. Artinya, model hampir tidak pernah salah menyebut anak normal sebagai stunting. Dalam praktik di Posyandu, hal ini berarti orang tua tidak akan menerima peringatan yang keliru, dan fasilitas kesehatan tidak akan dibebani rujukan yang sebenarnya tidak diperlukan. Kepercayaan masyarakat terhadap program skrining dapat terjaga.

Kenaikan *recall* menjadi 0.990 juga patut dicatat. Model mampu menemukan 99% kasus stunting yang sesungguhnya. Hanya satu dari setiap seratus anak stunting yang mungkin terlewat. Setiap anak yang terlewat memang sebuah kerugian, tetapi dengan sumber daya yang terbatas, mencapai *recall* setinggi ini adalah pencapaian yang sangat baik. Model berhasil menekan dua jenis kesalahan sekaligus, sebuah keseimbangan yang sulit dicapai dalam sistem klasifikasi pada umumnya.

AUC 1.000 yang dicapai oleh model tiga fitur memiliki makna penting. Metrik ini mengukur kemampuan model dalam membedakan kelas tanpa terikat pada satu ambang batas tertentu. Dengan *AUC* sempurna, model menempatkan seluruh balita stunting pada peringkat probabilitas yang lebih tinggi daripada seluruh balita normal. Tidak ada satu pun anak normal yang dianggap lebih berisiko daripada anak stunting. Kondisi ideal ini memberikan fleksibilitas penuh bagi petugas kesehatan untuk menyesuaikan ambang batas klasifikasi sesuai dengan ketersediaan sumber daya di lapangan, tanpa perlu khawatir urutan prioritasnya menjadi kacau.

Simpangan baku yang lebih kecil di seluruh metrik juga memberikan pesan yang jelas. Model tiga fitur tidak hanya lebih akurat, tetapi juga lebih stabil. Setiap variasi data yang dihadapkan kepadanya menghasilkan performa yang hampir seragam. Kestabilan ini merupakan indikator kuat bahwa model tidak bergantung pada karakteristik subset data tertentu, melainkan benar-benar menangkap pola umum yang membedakan anak stunting dari anak normal. Dalam konteks penerapan di dunia nyata, model yang stabil lebih dapat diandalkan karena ia akan tetap bekerja dengan baik meskipun data yang masuk memiliki karakteristik yang sedikit berbeda dari data pelatihan.

Peningkatan performa ini menegaskan bahwa keenam fitur yang dikeluarkan oleh algoritma Kneedle tidak memberikan informasi diskriminatif tambahan. Sebaliknya, kehadiran mereka justru menjadi derau yang mengganggu proses pencarian *hyperplane* optimal. Ketika derau tersebut dihilangkan, model dapat membangun margin pemisah yang lebih lebar dan lebih stabil (Hastie et al., 2009). Temuan ini sekaligus menjadi bukti empiris yang kuat bahwa proses seleksi fitur berbasis Kneedle telah mengambil keputusan yang tepat. Model dengan tiga fitur antropometri inti tidak hanya lebih sederhana, tetapi juga lebih akurat dan lebih andal. Kesederhanaan ini membawa implikasi praktis yang sangat berarti: petugas Posyandu hanya perlu mencatat umur, mengukur tinggi badan, dan mengetahui jenis kelamin anak untuk mendapatkan klasifikasi stunting yang sangat akurat. Tidak ada lagi kebutuhan untuk mengisi banyak formulir atau mencari data tambahan yang membebani. Inilah bentuk ideal dari sistem klasifikasi yang dapat

menjembatani antara kekuatan pembelajaran mesin dan realitas lapangan di Posyandu Indonesia.

4.1.5 Hasil Pengujian Konfigurasi Fitur Tidak Terpilih

Enam fitur yang tidak termasuk dalam himpunan PFI diuji secara terpisah sebagai konfigurasi mandiri. Fitur-fitur ini terdiri atas Berat, Pendidikan Ibu, Status Perawatan Makanan anak, Status Perawatan Kesehatan Anak, Menerima vitamin A, dan Imunisasi rotavirus. Pengujian ini bukan sekadar pelengkap rangkaian eksperimen, melainkan berfungsi sebagai kontrol negatif yang menguji validitas proses seleksi fitur. Jika algoritma Kneedle benar-benar telah memisahkan sinyal dari derau dengan tepat, maka model yang hanya mengandalkan keenam fitur yang dikeluarkan seharusnya menunjukkan performa yang jauh lebih rendah. Hipotesis inilah yang diuji melalui konfigurasi ini.

Tuning hyperparameter untuk konfigurasi ini dijalankan dengan prosedur hold-out validation yang sama. Hasil pencarian menunjukkan bahwa parameter terbaik adalah $C = 10$, $\gamma = 0.01$, dan $\beta = 0.7$. Nilai γ yang sangat kecil langsung memberikan petunjuk awal tentang karakteristik himpunan fitur ini. Fungsi Gaussian dengan lebar yang besar terpaksa digunakan agar model dapat menangkap sinyal yang tersisa, karena dalam ruang fitur yang miskin informasi, titik-titik data saling berjauhan dan hanya fungsi yang sangat lebar yang mampu menjangkau keterkaitan di antara mereka. Kondisi ini sudah mengisyaratkan bahwa kandungan informasi dalam keenam fitur ini sangat terbatas.

Setelah parameter terbaik diperoleh, model dievaluasi menggunakan validasi silang 10-fold. Rincian metrik pada setiap fold disajikan pada Tabel 4.6.

Tabel 4. 6 Hasil Pengujian K-fold Fitur Tidak Terpilih

Fold	Accuracy	Precision	Recall	F1-Score	AUC
1	0.7414	0.7608	0.8063	0.7828	0.7763
2	0.7756	0.7742	0.8739	0.8211	0.7893
3	0.7838	0.7995	0.8464	0.8223	0.8315
4	0.7673	0.7737	0.8424	0.8066	0.7920
5	0.7805	0.7923	0.8560	0.8229	0.8328
6	0.7640	0.7539	0.8546	0.8011	0.8162
7	0.7508	0.7620	0.8213	0.7906	0.8150
8	0.7756	0.7717	0.8571	0.8122	0.8205
9	0.7888	0.8000	0.8619	0.8298	0.8311
10	0.7640	0.7877	0.8370	0.8116	0.8084
Rata-rata ± SD	0.769 ± 0.015	0.778 ± 0.016	0.846 ± 0.020	0.810 ± 0.015	0.811 ± 0.020

Hasil evaluasi menunjukkan penurunan performa yang sangat tajam. Akurasi hanya mencapai 0.769, yang berarti lebih dari 23% sampel diklasifikasikan secara keliru. *F1-score* terpuruk di angka 0.810, sementara *AUC* hanya 0.811. Jika dibandingkan dengan model tiga fitur PFI yang mencapai *F1-score* 0.994 dan *AUC* sempurna 1.000, jurang perbedaan ini sangat lebar. Selisih *F1-score* sebesar 0.184 bukanlah selisih yang dapat diabaikan; ini adalah bukti kuat bahwa keenam fitur ini hampir tidak memiliki daya diskriminasi yang memadai.

Pola metrik yang muncul juga menarik untuk dicermati. *Recall* rata-rata mencapai 0.846, sementara presisi hanya 0.778. Kesenjangan antara kedua metrik ini sangat lebar. Model masih mampu menemukan sekitar 84,6% kasus stunting yang sebenarnya, tetapi harus membayarnya dengan sangat mahal: lebih dari 22% prediksi stunting yang dihasilkan adalah keliru. Model bekerja dengan cara yang longgar, mengklasifikasikan banyak anak sebagai stunting untuk menghindari kehilangan kasus. Strategi ini mungkin terlihat menyelamatkan banyak anak, tetapi dampak buruknya justru lebih besar. Positif palsu dalam jumlah besar akan

membanjiri fasilitas kesehatan dengan rujukan yang tidak perlu, menguras sumber daya yang seharusnya bisa dialokasikan untuk anak yang benar-benar membutuhkan. Lebih dari itu, orang tua yang menerima peringatan keliru akan mengalami kecemasan, dan jika hal ini terjadi berulang, kepercayaan mereka terhadap program skrining akan terkikis.

Simpangan baku yang besar di semua metrik semakin mempertegas ketidakstabilan model. Simpangan baku *recall* mencapai 0.020, menunjukkan bahwa kemampuan model dalam menemukan kasus stunting sangat bergantung pada komposisi data yang kebetulan muncul di setiap lipatan. Pada beberapa lipatan, model sedikit lebih beruntung; pada lipatan lain, kemampuannya menurun drastis. Ini adalah ciri khas model yang bekerja di ruang fitur yang hampir tidak mengandung informasi. Model tidak memiliki cukup sinyal untuk membangun batas keputusan yang konsisten, sehingga performanya naik-turun secara tidak terkendali.

Nilai gamma optimal sebesar 0.01 yang terpilih semakin mengonfirmasi rendahnya kandungan informasi. Fungsi *Gaussian* yang sangat lebar adalah strategi terakhir model untuk menangkap sisa-sisa sinyal yang nyaris tidak ada. Dalam ruang fitur yang kaya informasi, model biasanya memerlukan fungsi yang lebih sempit untuk menangkap perbedaan-perbedaan halus antarkasus. Sebaliknya, dalam ruang fitur yang miskin, model terpaksa memperlebar pengaruh setiap titik data agar dapat menemukan keterkaitan apa pun yang masih tersisa. Model tidak lagi membangun batas keputusan yang tajam; ia bekerja dalam mode menebak secara kasar.

Temuan ini menutup seluruh rangkaian pengujian dengan simpulan yang kokoh. Algoritma Kneedle telah mengambil keputusan yang tepat ketika menempatkan titik siku pada fitur Jenis Kelamin. Keenam fitur yang dikeluarkan, yaitu Berat, Pendidikan Ibu, Status Perawatan Makanan anak, Status Perawatan Kesehatan Anak, Menerima vitamin A, dan Imunisasi rotavirus, terbukti tidak memiliki kapasitas diskriminatif yang memadai. Informasi yang mereka bawa terlalu lemah dan tercampur derau, sehingga bukan hanya tidak membantu, tetapi justru mengganggu jika dipaksakan masuk ke dalam model. Fakta bahwa model dengan tiga fitur antropometri inti mampu mencapai performa nyaris sempurna, sementara model dengan enam fitur lainnya justru terpuruk, menjadi bukti empiris yang sulit dibantah. Proses seleksi fitur yang dilakukan bukanlah pemangkasan sembarangan, melainkan pemisahan yang presisi antara sinyal dan derau. Inilah validasi paling kuat bahwa titik siku yang terdeteksi oleh algoritma Kneedle memang menjadi batas alami antara fitur yang informatif dan fitur yang tidak.

4.1.6 Hasil Perbandingan Kinerja 3 Konfigurasi

Setelah seluruh konfigurasi dievaluasi dengan prosedur yang sama, perbandingan performa secara langsung dapat dilakukan. Ringkasan metrik dari ketiga konfigurasi disajikan pada Tabel 4.7.

Tabel 4. 7 Perbandingan Kinerja Model pada Tiga Konfigurasi Fitur

Metrik Evaluasi	9 Fitur	3 Fitur PFI	Fitur Tidak Terpilih
Accuracy	0.978 ± 0.007	0.993 ± 0.003	0.769 ± 0.015
Precision	0.984 ± 0.007	0.997 ± 0.002	0.778 ± 0.016
Recall	0.979 ± 0.008	0.990 ± 0.005	0.846 ± 0.020
F1-Score	0.981 ± 0.006	0.994 ± 0.002	0.810 ± 0.015
AUC	0.998 ± 0.001	1.000 ± 0.000	0.811 ± 0.020

Perbandingan pada Tabel 4.7 menegaskan beberapa temuan penting yang saling mengonfirmasi. Model dengan tiga fitur PFI tidak hanya mengungguli model fitur tidak terpilih, tetapi juga melampaui model sembilan fitur di semua metrik. Akurasi model PFI mencapai 0.993, lebih tinggi 0.015 dari model sembilan fitur dan lebih tinggi 0.224 dari model fitur tidak terpilih. Selisih akurasi antara model PFI dan model fitur tidak terpilih sangat besar, menunjukkan bahwa keenam fitur yang dikeluarkan hampir tidak memiliki kemampuan untuk membedakan kelas stunting dan normal secara akurat.

Kesenjangan paling mencolok terlihat pada *F1-score*. Model PFI mencatat *F1-score* 0.994, sementara model fitur tidak terpilih hanya mencapai 0.810. Selisih sebesar 0.184 ini sangat signifikan secara statistik dan praktis. Fakta bahwa model dengan tiga fitur mampu melampaui model dengan sembilan fitur membuktikan bahwa keenam fitur tambahan bukan hanya tidak membantu, melainkan justru menjadi pengganggu yang menurunkan performa. Kehadiran mereka dalam model sembilan fitur menambahkan derau yang mengaburkan batas keputusan, sehingga *F1-score* model sembilan fitur tertahan di 0.981, lebih rendah 0.013 dari model PFI.

Recall yang tinggi pada model fitur tidak terpilih, yaitu 0.846, tidak boleh ditafsirkan sebagai kemampuan yang baik. Angka ini dicapai dengan mengorbankan presisi yang sangat rendah. Lebih dari 22% prediksi stunting yang dihasilkan oleh model fitur tidak terpilih adalah keliru. Kombinasi *recall* tinggi dan presisi rendah adalah pola khas model yang bekerja secara longgar, mengklasifikasikan hampir semua kasus sebagai positif demi menghindari

kehilangan kasus stunting. Strategi ini tidak sesuai untuk skrining kesehatan karena akan membanjiri sistem rujukan dengan kasus yang sebenarnya tidak memerlukan penanganan.

AUC memberikan perspektif yang paling tajam tentang kualitas diskriminasi masing-masing konfigurasi. Model PFI mencapai *AUC* sempurna 1.000, yang berarti model mampu menempatkan seluruh balita stunting pada peringkat probabilitas yang lebih tinggi daripada seluruh balita normal tanpa satu pun kesalahan peringkat. Model sembilan fitur juga sangat baik dengan *AUC* 0.998. Sementara itu, model fitur tidak terpilih hanya mencapai *AUC* 0.811, yang hampir mendekati batas bawah kemampuan diskriminasi yang dapat diterima. Simpangan baku *AUC* pada model fitur tidak terpilih yang mencapai 0.020 juga menandakan ketidakstabilan yang tinggi. Pada beberapa lipatan, model sedikit lebih beruntung; pada lipatan lain, kemampuannya menurun drastis.

Simpangan baku di seluruh metrik juga menceritakan kisah yang konsisten. Model PFI memiliki simpangan baku paling kecil, menandakan kestabilan tertinggi. Model sembilan fitur memiliki simpangan baku yang sedikit lebih besar, menunjukkan bahwa kehadiran fitur-fitur yang kurang informatif sedikit meningkatkan variabilitas performa. Model fitur tidak terpilih memiliki simpangan baku terbesar, menegaskan bahwa model sangat tidak stabil ketika hanya mengandalkan fitur-fitur yang minim informasi. Kestabilan yang rendah adalah masalah serius dalam konteks penerapan di dunia nyata, karena model harus dapat diandalkan dalam berbagai kondisi data yang mungkin berbeda dari data pelatihan.

Secara keseluruhan, perbandingan ketiga konfigurasi ini memberikan bukti yang kuat dan konsisten. Algoritma Kneedle telah mengambil keputusan yang tepat dengan menempatkan titik siku pada fitur Jenis Kelamin. Tiga fitur antropometri inti, yaitu Umur Per Bulan, Tinggi, dan Jenis Kelamin, sudah cukup untuk membangun model klasifikasi stunting yang sangat akurat, stabil, dan efisien. Keenam fitur lainnya bukan hanya tidak diperlukan, tetapi justru menjadi derau yang mengganggu. Temuan ini menutup seluruh rangkaian pengujian dengan simpulan yang kokoh: proses seleksi fitur yang dilakukan bukanlah pemangkasan sembarangan, melainkan pemisahan yang presisi antara sinyal dan derau.

4.2 Pembahasan

Model SVM-MKL dengan sembilan fitur telah menunjukkan performa yang baik dalam mengklasifikasikan status stunting. Proses optimasi hyperparameter melalui hold-out validation berhasil menemukan konfigurasi yang tidak hanya unggul dalam metrik keseluruhan, tetapi juga menjaga keseimbangan yang sehat antara presisi dan recall. Pemilihan parameter terbaik tidak bertumpu semata-mata pada F1-score tertinggi. Kombinasi ke-22 dengan $\beta = 0.5$ memang mencatat F1-score paling tinggi, tetapi kombinasi ke-23 dengan $\beta = 0.6$ dipilih karena memberikan presisi yang lebih baik dengan pengorbanan F1-score yang sangat kecil. Keputusan ini mencerminkan pemahaman bahwa dalam konteks skrining stunting, kesalahan positif palsu dan negatif palsu sama-sama membawa konsekuensi serius. Positif palsu dapat menimbulkan kecemasan orang tua dan membebani fasilitas kesehatan dengan rujukan yang tidak perlu. Negatif palsu dapat menghilangkan kesempatan anak untuk mendapatkan intervensi gizi pada masa

pertumbuhan kritis yang tidak dapat diulang. Model dengan $\beta = 0.6$ memberikan presisi yang lebih tinggi, sehingga lebih sedikit anak normal yang dirujuk sebagai stunting. Properti ini penting untuk menjaga kepercayaan masyarakat terhadap program skrining.

Evaluasi menggunakan validasi silang memperlihatkan bahwa model memiliki kestabilan yang tinggi. Simpangan baku yang kecil di semua metrik menandakan bahwa model tidak bergantung pada komposisi data tertentu. Model mampu mempertahankan konsistensinya di seluruh variasi subset data. Bahkan pada lipatan dengan performa terendah, kemampuan diskriminasi model tetap bertahan pada tingkat yang sangat tinggi. AUC pada fold terlemah masih mencatat 0.9975, yang menunjukkan bahwa penurunan akurasi lebih disebabkan oleh pergeseran titik potong klasifikasi, bukan oleh melemahnya kemampuan model dalam mengurutkan risiko. Stabilitas ini menunjukkan bahwa margin pemisah yang dibangun oleh SVM-MKL benar-benar lebar dan tidak mudah terganggu oleh perbedaan karakteristik data antar lipatan. Kemampuan generalisasi semacam ini adalah prasyarat penting sebelum model diterapkan pada populasi yang lebih luas.

Pendekatan MKL dengan bobot $\beta = 0.6$ memberikan kontribusi penting terhadap performa model. Proporsi ini menunjukkan bahwa hubungan linear antara fitur antropometri dan status stunting sedikit lebih dominan dibandingkan interaksi non-linear. Tinggi badan yang rendah terhadap umur berkorelasi langsung dengan kondisi stunting, dan hubungan linear ini ditangkap dengan baik oleh kernel linear. Meskipun demikian, kontribusi kernel RBF yang tetap cukup besar mengindikasikan bahwa interaksi non-linear antarfaktor tetap relevan. Kombinasi

beberapa faktor risiko dapat saling memperkuat secara tidak proporsional, dan kernel RBF hadir untuk menangkap pola-pola semacam ini. Kedua jenis hubungan tersebut sama-sama eksis dalam data, dan MKL memungkinkan model memanfaatkan keduanya secara bersamaan.

Hasil Permutation Feature Importance membuka wawasan yang sangat berharga tentang struktur informasi dalam data stunting. Fitur Umur Per Bulan mencatat penurunan kemampuan prediksi yang sangat dominan. Temuan ini menegaskan bahwa lebih dari setengah kemampuan model bertumpu pada informasi umur. Hal ini sangat masuk akal secara klinis. Stunting didefinisikan sebagai tinggi badan yang rendah terhadap umur, sehingga umur menjadi sumbu utama yang menentukan posisi seorang anak dalam kurva pertumbuhan. Tanpa informasi umur, tinggi badan kehilangan makna diagnostiknya. Fitur Tinggi dan Jenis Kelamin melengkapi tiga besar fitur yang paling berkontribusi. Ketiga fitur ini adalah komponen utama dalam penghitungan *Height-for-Age Z-score* (HAZ), standar diagnosis stunting yang ditetapkan oleh WHO (WHO Multicentre Growth Reference Study Group, 2006). Fakta bahwa algoritma Kneedle menempatkan titik siku tepat pada fitur ketiga menunjukkan bahwa data itu sendiri yang mengarahkan proses seleksi. Algoritma tidak membutuhkan asumsi subjektif. Pola penurunan yang melandai tajam setelah Jenis Kelamin sudah cukup untuk menandai batas antara sinyal dan derau.

Pengujian pada subset fitur memberikan bukti yang paling meyakinkan. Model yang hanya menggunakan tiga fitur antropometri justru menunjukkan peningkatan performa di semua metrik dibandingkan model dengan sembilan fitur.

Peningkatan ini menegaskan bahwa fitur-fitur yang dikeluarkan tidak memberikan informasi diskriminatif tambahan. Sebaliknya, kehadiran mereka justru mengganggu proses pencarian hyperplane optimal. Fenomena ini konsisten dengan prinsip curse of dimensionality. Fitur-fitur yang tidak membawa informasi hanya akan menambah derau dan mempersulit algoritma berbasis jarak seperti SVM dalam membangun margin pemisah yang optimal (Hastie et al., 2009).

Pergeseran nilai gamma optimal dari 0.1 menjadi 1.0 pada model subset fitur juga layak dicermati. Ketika ruang fitur mengecil dari sembilan dimensi menjadi tiga dimensi, data menjadi lebih padat dan jarak antartitik menjadi lebih bermakna. Model memerlukan fungsi Gaussian yang lebih sempit agar dapat menangkap perbedaan halus antarkasus dalam ruang yang lebih sederhana. Ini adalah respons adaptif model terhadap penyederhanaan ruang fitur. Perilaku ini menunjukkan bahwa model tidak kaku terhadap perubahan dimensi, melainkan mampu menyesuaikan kompleksitasnya sesuai dengan karakteristik data yang dihadapi.

Pengujian pada fitur tidak terpilih menjadi kontrol negatif yang memperkuat seluruh temuan. Model yang hanya menggunakan fitur-fitur yang dikeluarkan oleh Kneedle menunjukkan performa yang sangat rendah. Kemampuan diskriminasinya tertinggal jauh dibandingkan model PFI. Selisih yang sangat besar ini tidak mungkin diabaikan dan secara statistik sangat signifikan. Recall yang cukup tinggi pada model ini tidak boleh disalahartikan sebagai kemampuan yang baik. Angka tersebut dicapai dengan mengorbankan presisi yang sangat rendah, yang berarti sebagian besar prediksi stunting yang dihasilkan model adalah keliru. Pola recall

tinggi dan presisi rendah adalah ciri khas model yang bekerja secara longgar, mengklasifikasikan hampir semua kasus sebagai positif untuk menghindari kehilangan kasus stunting. Strategi ini tidak sesuai untuk skrining kesehatan karena akan membanjiri sistem rujukan dengan kasus yang sebenarnya tidak memerlukan penanganan.

Nilai gamma optimal yang sangat kecil pada konfigurasi fitur tidak terpilih semakin mengonfirmasi rendahnya kandungan informasi dalam fitur-fitur tersebut. Model terpaksa menggunakan fungsi Gaussian yang sangat lebar untuk menangkap sisa-sisa sinyal yang hampir tidak ada. Model bekerja dalam mode menebak secara kasar karena tidak terdapat cukup informasi untuk membangun batas keputusan yang tajam. Semua bukti ini mengarah pada satu kesimpulan yang kokoh. Algoritma Kneedle telah mengambil keputusan yang tepat. Proses seleksi fitur yang dilakukan bukanlah pemangkasan sembarangan, melainkan pemisahan yang presisi antara sinyal dan derau.

Secara praktis, temuan ini membawa implikasi yang menggembirakan. Petugas Posyandu hanya perlu mencatat umur, mengukur tinggi badan, dan mengetahui jenis kelamin anak untuk memperoleh klasifikasi stunting yang sangat andal. Ketiga data ini sudah menjadi bagian dari rutinitas penimbangan bulanan. Tidak ada kebutuhan untuk mengumpulkan data tambahan yang dapat membebani petugas. Kesederhanaan ini memperbesar peluang adopsi sistem di lapangan, termasuk di wilayah dengan akses data yang terbatas. Model dengan tiga fitur tidak hanya lebih akurat, tetapi juga lebih efisien dan lebih mudah diimplementasikan.

Inilah bentuk ideal dari sistem klasifikasi yang dapat menjembatani antara kekuatan pembelajaran mesin dan realitas lapangan di Posyandu Indonesia.

4.3 Pembahasan Integrasi Islam

Hasil pengujian dan pembahasan yang telah diuraikan sebelumnya tidak hanya berbicara tentang angka dan performa model. Pembahasan ini juga menyentuh persoalan yang lebih luas, yaitu bagaimana memastikan setiap anak Indonesia tumbuh dengan sehat dan memiliki kesempatan yang sama untuk berkembang. Dalam konteks inilah nilai-nilai Islam menemukan relevansinya, bukan sebagai tempelan, melainkan sebagai landasan yang mendorong lahirnya kepedulian dan tindakan nyata.

Al-Qur'an telah mengingatkan tentang bahaya meninggalkan generasi yang lemah. Allah Subhanahu wa Ta'ala berfirman dalam QS. An-Nisa' ayat 9

وَلْيَحْشَ الَّذِينَ لَوْ تَرَكُوا مِنْ خَلْفِهِمْ ذُرِّيَّةً ضِعْفًا خَافُوا عَلَيْهِمْ فَلْيَتَّقُوا اللَّهَ وَلْيَقُولُوا قَوْلًا سَدِيدًا

“Dan hendaklah takut (kepada Allah) orang-orang yang sekiranya mereka meninggalkan keturunan yang lemah di belakang mereka yang mereka khawatir terhadap (kesejahteraan)nya. Oleh sebab itu, hendaklah mereka bertakwa kepada Allah dan berbicara dengan tutur kata yang benar.” (QS. An-Nisa’/4:9).

Menurut M. Quraish Shihab dalam Tafsir Al-Mishbah, ayat ini ditujukan kepada orang-orang yang berada di sekeliling pemilik harta yang sedang sakit parah. Mereka sering kali memberi nasihat agar yang sakit itu mewasiatkan sebagian hartanya kepada orang lain, sehingga anak-anaknya sendiri terbengkalai. Ayat ini memperingatkan mereka: hendaklah mereka membayangkan seandainya mereka sendiri yang akan meninggalkan anak-anak yang lemah di belakang. Apakah mereka akan menerima nasihat seperti yang mereka berikan. Tentu tidak.

Oleh karena itu, hendaklah mereka takut kepada Allah dan mengucapkan perkataan yang benar lagi tepat, qaulan sadidan. Kata sadidan tidak hanya berarti benar, tetapi juga tepat sasaran, serta mengandung makna meruntuhkan sesuatu kemudian memperbaikinya, yakni kritik yang membangun. Ayat ini juga mengajarkan bahwa setiap muslim harus khawatir meninggalkan generasi yang lemah secara fisik, mental, dan spiritual. Hendaklah ia bertindak preventif, termasuk dalam hal pemenuhan gizi anak untuk mencegah kondisi seperti stunting.

Stunting adalah salah satu bentuk kelemahan yang paling nyata. Anak yang mengalaminya tidak hanya tumbuh lebih pendek, tetapi juga menghadapi hambatan dalam perkembangan otak yang akan memengaruhi seluruh hidupnya. Sistem klasifikasi yang dibangun dalam penelitian ini lahir dari kesadaran bahwa mencegah lebih baik daripada mengobati. Setiap anak berhak mendapatkan perhatian sebelum terlambat.

Tanggung jawab untuk memuliakan anak juga ditegaskan dalam sabda Rasulullah ﷺ yang diriwayatkan oleh sahabat Anas bin Malik radhiyallahu anhu

وَقَالَ عَلَيْهِ الصَّلَاةُ وَالسَّلَامُ: { أَكْرِمُوا أَوْلَادَكُمْ وَأَحْسِنُوا آدَابَهُمْ

Nabi ﷺ bersabda, "Muliaikanlah anak-anak kalian dan ajarilah mereka tata krama." (HR. Ibnu Majah)

Hadits ini memberikan kerangka berpikir yang sangat relevan dengan upaya pencegahan dan penanganan stunting. Konsep mukmin yang kuat tidak terbatas pada kekuatan fisik atau iman semata, tetapi juga mencakup kekuatan ilmu, teknologi, dan sistem yang memungkinkan deteksi dini serta intervensi yang tepat sasaran. Penerapan metode Support Vector Machine dengan pendekatan Multiple Kernel Learning merupakan perwujudan nyata dari perintah untuk bersungguh

sungguh mendapatkan apa yang bermanfaat. Sistem klasifikasi stunting berbasis web yang dikembangkan adalah alat bantu agar tenaga kesehatan maupun orang tua dapat mendeteksi risiko stunting secara lebih akurat dan cepat. Ini adalah bentuk ikhtiar ilmiah yang tidak bertentangan dengan tawakal. Justru sejalan dengan perintah untuk menggunakan akal dan teknologi sebagai karunia Allah.

Pada akhirnya, seluruh proses yang telah dilalui dalam penelitian ini adalah cerminan dari semangat ikhtiar yang diajarkan dalam Islam. Allah SWT berfirman dalam QS. An-Najm ayat 39

وَأَنْ لَّيْسَ لِلْإِنْسَانِ إِلَّا مَا سَعَى

“Dan bahwa manusia hanya memperoleh apa yang telah diusahakannya.” (QS. An-Najm/53:39).

M. Quraish Shihab dalam Tafsir Al-Mishbah menjelaskan bahwa kata sa'a yang berarti berusaha pada mulanya berarti berjalan cepat, kemudian digunakan dalam arti berupaya secara sungguh-sungguh. Ayat ini menegaskan bahwa setiap manusia hanya akan mendapatkan balasan sesuai dengan amal usahanya sendiri. Seseorang tidak dapat memikul dosa orang lain. Sebaliknya, ia juga tidak dapat meraih manfaat dari amal orang lain tanpa usahanya sendiri. Dalam konteks perawatan anak dan stunting, ayat ini menjadi pendorong yang sangat kuat. Orang tua tidak bisa hanya berharap anaknya tumbuh sehat tanpa melakukan usaha yang nyata. Kesehatan anak tidak datang dengan sendirinya. Ia memerlukan usaha sungguh-sungguh, mulai dari pemberian gizi yang cukup, pemantauan pertumbuhan rutin di Posyandu, hingga pemanfaatan alat bantu seperti sistem klasifikasi yang dikembangkan dalam penelitian ini. Setiap langkah kecil yang

diambil oleh orang tua, setiap kunjungan ke Posyandu, setiap perhatian terhadap tinggi dan berat badan anak, semuanya adalah wujud sa'a yang akan diperhitungkan. Peneliti menyadari sepenuhnya bahwa tanpa pertolongan Allah SWT, semua usaha ini tidak akan mencapai hasil yang diharapkan. Semoga sistem yang dihasilkan dapat menjadi bagian dari ikhtiar bersama menjaga generasi penerus bangsa.

Penelitian ini pada akhirnya bukan sekadar tentang akurasi model atau metrik evaluasi. Ia adalah tentang bagaimana ilmu pengetahuan diarahkan untuk menjawab panggilan moral yang telah digariskan dalam Al-Qur'an dan Sunnah. Sistem klasifikasi stunting yang dibangun tidak dimaksudkan untuk menggantikan peran orang tua atau tenaga kesehatan, melainkan untuk memperkuat mereka dengan informasi yang tepat dan dapat diandalkan. Ketika teknologi dimanfaatkan untuk memastikan tidak ada anak yang tertinggal dalam pertumbuhannya, maka ia telah menjalankan fungsi kekhalifahan yang diamanatkan kepada manusia: menjaga, memelihara, dan memuliakan kehidupan. Inilah makna terdalam dari integrasi Islam dalam penelitian ini. Bukan sekadar penyisipan ayat dan hadits, melainkan upaya untuk menjadikan nilai-nilai tersebut sebagai arah dan tujuan dari setiap langkah ilmiah yang diambil.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Model Support Vector Machine dengan pendekatan Multiple Kernel Learning berhasil diimplementasikan untuk klasifikasi stunting pada anak balita. Konfigurasi terbaik diperoleh melalui optimasi hyperparameter yang mempertimbangkan keseimbangan presisi dan recall, dan model telah diintegrasikan ke dalam aplikasi berbasis web untuk digunakan di Posyandu.

Model dengan sembilan fitur menunjukkan kinerja tinggi dan stabil di seluruh metrik evaluasi, dengan simpangan baku kecil yang menandakan kemampuan generalisasi yang baik. Keseimbangan antara presisi dan recall menjadikan model ini andal dalam menekan kesalahan positif palsu dan negatif palsu secara bersamaan.

Analisis Permutation Feature Importance yang diproses melalui algoritma Kneedle hanya mempertahankan tiga fitur antropometri inti, yaitu Umur Per Bulan, Tinggi, dan Jenis Kelamin. Pengurangan enam fitur lainnya justru meningkatkan akurasi menjadi 0.993, F1-score menjadi 0.994, dan AUC mencapai 1.000. Pengujian pada fitur tidak terpilih membuktikan bahwa fitur-fitur tersebut hanya mencapai F1-score 0.810 dan AUC 0.811, menegaskan bahwa variabel antropometri dasar sudah sangat kuat sebagai prediktor stunting. Model dengan sembilan fitur tetap diintegrasikan ke dalam sistem berbasis web karena

menyediakan informasi kontekstual yang lebih lengkap bagi tenaga kesehatan di lapangan.

5.2 Saran

Pengembangan sistem ini masih dapat dilanjutkan melalui beberapa langkah. Data primer yang diambil langsung dari Posyandu di berbagai daerah akan membuat model lebih mencerminkan kondisi terkini di lapangan. Validasi dengan data primer juga penting untuk memastikan bahwa performa model yang tinggi tidak hanya berlaku pada dataset IFLS-5, tetapi juga pada populasi balita Indonesia saat ini. Model juga dapat dikembangkan ke arah klasifikasi multi-kelas, misalnya dengan menambahkan kategori *severely stunting*, sehingga sistem mampu memberikan informasi status gizi yang lebih terperinci. Eksplorasi kernel lain seperti *polynomial* atau *sigmoid* dalam kerangka MKL layak dipertimbangkan untuk menguji potensi peningkatan performa klasifikasi.

Di sisi implementasi, sistem berbasis web yang telah dibangun dapat dilengkapi dengan fitur visualisasi pertumbuhan anak dan sistem peringatan dini berbasis tren pengukuran. Integrasi dengan sistem informasi kesehatan nasional akan mempermudah sinkronisasi data di tingkat puskesmas dan dinas kesehatan. Pengujian penerimaan pengguna terhadap aplikasi web juga sangat disarankan. Pengujian ini akan memastikan bahwa sistem mudah digunakan dan sesuai dengan kebutuhan tenaga kesehatan di lapangan.

DAFTAR PUSTAKA

- Adzhima, F. (2023). Klasifikasi Status Stunting Balita Dengan Metode Support Vector Machine Berbasis Web. *INOVTEK Polbeng - Seri Informatika*, 8(2), 381. <https://doi.org/10.35314/isi.v8i2.3641>
- Andi, T., Ismail, A. R., Pranolo, A., & Kusuma, C. J. C. (2026). Classification of Heart Disease with Machine Learning: A Comparison of Grid Search, Random Search, and Bayesian Optimization. *International Journal on Informatics Visualization*, 10(1), 262–270. <http://irep.iium.edu.my/127509/>
- Angga Aditya Permana, M. F. R. (2024). Implementasi Algoritma Support Vector Machine Untuk Klasifikasi Status Stunting Pada Balita. *G-Tech : Jurnal Teknologi Terapan*, 8(1), 186–195. <https://ejournal.uniramalang.ac.id/index.php/g-tech/article/view/1823/1229>
- Anggraini, Y., & Rachmawati, Y. (2021). Preventing Stunting in Children. *Proceedings of the 5th International Conference on Early Childhood Education (ICECE 2020)*, 538(Icece 2020), 203–206. <https://doi.org/10.2991/assehr.k.210322.044>
- Aprihartha, M. A., & Idham, I. (2024). Optimization of Classification Algorithms Performance with k-Fold Cross Validation. *Eigen Mathematics Journal*, 7(2), 61–66. <https://doi.org/10.29303/emj.v7i2.212>
- Bitew, F. H., Sparks, C. S., & Nyarko, S. H. (2025). Predicting Stunting Status among Under Five Children in Ethiopia Using Ensemble Machine Learning Algorithms. *Scientific Reports*.
- Bradley, A. P. (1997). The Use of the Area under the ROC Curve in the Evaluation of Machine Learning Algorithms. *Pattern Recognition*, 30(7), 1145–1159. [https://doi.org/10.1016/S0031-3203\(96\)00142-2](https://doi.org/10.1016/S0031-3203(96)00142-2)
- Chicco, D., & Jurman, G. (2020). The Advantages of the Matthews Correlation Coefficient MCC over 1F Score and Accuracy in Binary Classification Evaluation. *BMC Genomics*, 21(1), 6. <https://doi.org/10.1186/s12864-019-6413-7>
- Cortes, C., & Vapnik, V. (1995). Support-Vector Networks. *Machine Learning*, 20(3), 273–297. <https://doi.org/10.1007/BF00994018>
- Fardila elba, Hassan, H. C., Umar, N. S., & Hilmanto, D. (2023). Stunting Interventions in Developing Countries: Literature Review. *International Journal of Health Sciences*, 1(3), 408–418. <https://doi.org/10.59585/ijhs.v1i3.146>
- Fawcett, T. (2006). *An introduction to ROC analysis*. 27, 861–874. <https://doi.org/10.1016/j.patrec.2005.10.010>

- Gaffar, A. W. M., Halis, A. M., Purnawansyah, P., & Jabir, S. R. (2024). Penerapan Algoritma Support Vector Machine untuk Klasifikasi Stunting pada Balita di Kabupaten Enrekang. *Jurnal Minfo Polgan*, 13(1), 286–292. <https://doi.org/10.33395/jmp.v13i1.13620>
- Gönen, M., & Alpaydin, E. (2011). Multiple Kernel Learning Algorithms. *Journal of Machine Learning Research*, 12, 2211–2268. <https://dl.acm.org/doi/10.5555/1953048.2021035>
- Guyon, I., & Elisseeff, A. (2003). An Introduction to Variable and Feature Selection. *Journal of Machine Learning Research*, 3, 1157–1182. <https://www.jmlr.org/papers/volume3/guyon03a/guyon03a.pdf>
- Hasdyna, N., Dinata, R. K., Rahmi, & Fajri, T. I. (2024). Hybrid Machine Learning for Stunting Prevalence: A Novel Comprehensive Approach to Its Classification, Prediction, and Clustering Optimization in Aceh, Indonesia. *Informatics*, 11(4). <https://doi.org/10.3390/informatics11040089>
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The Elements of Statistical Learning: Data Mining, Inference, and Prediction* (2nd ed.). Springer.
- Hicks, S. A., Strümke, I., Thambawita, V., Hammou, M., Riegler, M. A., Halvorsen, P., & Parasa, S. (2022). On Evaluation Metrics for Medical Applications of Artificial Intelligence. *Scientific Reports*, 12(1), 5979. <https://doi.org/10.1038/s41598-022-09954-8>
- Hidayati, S. F., Umboh, V., & Rondonuwu, S. H. E. (2022). Relationship between Nutritional Status and Urinary Tract Infection in Children. *E-CliniC*, 10(2), 288. <https://doi.org/10.35790/ecl.v10i2.37830>
- Hughes, G. F. (1968). On the Mean Accuracy of Statistical Pattern Recognizers. *IEEE Transactions on Information Theory*, 14(1), 55–63. <https://doi.org/10.1109/TIT.1968.1054102>
- Kemenkes RI. (2022). Hasil Survei Status Gizi Indonesia (SSGI) 2022. *Kemenkes*, 1–150.
- Kementerian PPN. (2025). *Rencana Pembangunan Jangka Menengah Nasional*. 1–23.
- Nghien, N. B., Cong, C. N., Thi, N. N., & Dung, V. Q. (2025). Comparison among Search Algorithms for Hyperparameter of Support Vector Machine Optimization. *IAES International Journal of Artificial Intelligence (IJ-AI)*, 14(5), 3802–3815. <https://doi.org/10.11591/ijai.v14.i5.pp3802-3815>
- Rahmi, I., Susanti, M., Yozza, H., & Wulandari, F. (2022). Classification of Stunting in Children Under Five Years in Padang City Using Support Vector Machine. *Barekeng*, 16(3), 771–778. <https://doi.org/10.30598/barekengvol16iss3pp771-778>
- Schölkopf, B., & Smola, A. J. (Eds.). (2001). Designing Kernels. In *Learning with*

Kernels: Support Vector Machines, Regularization, Optimization, and Beyond (p. 0). The MIT Press. <https://doi.org/10.7551/mitpress/4175.003.0018>

- Somol, P., Novovičová, J., & Pudil, P. (2010). Improving Feature Selection Process Resistance to Failures Caused by Curse-of-Dimensionality Effects. *Kybernetika*, 46(3), 401–425. <http://www.kybernetika.cz/content/2010/3/401>
- Sugiarti, S., Sunarsih, S., Kohir, D. S., Rahmayati, E., & Purwati, P. (2023). Efektivitas program positive deviance terhadap peningkatan status gizi balita melalui kegiatan pos gizi: Literature review. *THE JOURNAL OF Mother and Child Health Concerns*, 3(1), 9–20. <https://doi.org/10.56922/mchc.v3i1.336>
- Tougui, I., Jilbab, A., & Mhamdi, J. El. (2021). Impact of the Choice of Cross-Validation Techniques on the Results of Machine Learning-Based Diagnostic Applications. *Healthcare Informatics Research*, 27(3), 189–199. <https://doi.org/10.4258/HIR.2021.27.3.189>
- UNICEF. (2021). Levels and trends in child malnutrition UNICEF / WHO / World Bank Group Joint Child Malnutrition Estimates Key findings of the 2021 edition. *World Health Organization*, 1–32. <https://www.who.int/publications/i/item/9789240025257>
- Yunita, A., Fatichah, C., Yuhana, U. L., Informatika, J. T., Islam, U., Malik, M., Malang, I., Teknik, J., Fakultas, I., & Informasi, T. (2012). IMPLEMENTASI METODE MULTIPLE KERNEL SUPPORT VECTOR MACHINE UNTUK SELEKSI FITUR DARI DATA EKSPRESI GEN DENGAN STUDI KASUS LEUKIMIA DAN TUMOR USUS BESAR. *Matics Jurnal Ilmu Komputer Dan Teknologi Informasi (Journal of Computer Science and Information Technology)*. <https://doi.org/10.18860/mat.v0i0.1563>