

KLASIFIKASI AUTISME PADA ANAK BERDASARKAN *DIAGNOSTIC AND STATISTICAL MANUAL OF MENTAL DISORDERS* (DSM-5) MENGGUNAKAN *SUPPORT VECTOR MACHINE* (SVM)

Avinda Nur Azizah¹,
Siti Umami Masruroh², Neneng Tati Sumiati³

^{1,2}Department of Informatics Engineering, Faculty of Science and Technology, State Islamic University Syarif Hidayatullah Jakarta

^{1,2}Jl.Ir H. Juanda Street Number 95, Ciputat, West Tangerang, Banten, Indonesia
Email: avinda.nurazizah21@mhs.uinjkt.ac.id¹, ummi.masruroh@uinjkt.ac.id²,
neneng.tati@uinjkt.ac.id³

ABSTRACT

Article:

Accepted: xxx xx, 20xx

Revised: xxx xx, 20xx

Issued: xxx xx, 20xx

***Correspondence Address:**
(email)

Autisme adalah gangguan perkembangan yang memengaruhi kemampuan komunikasi dan interaksi sosial anak. Diagnosis dini sangat penting untuk mencegah keterlambatan dalam penanganan, namun kurangnya kesadaran orang tua terhadap gejala awal, serta belum tersedianya alat bantu klasifikasi untuk mengidentifikasi tingkat keparahan autisme berdasarkan kriteria DSM-5. Penelitian ini bertujuan mengklasifikasikan autisme pada anak berdasarkan kriteria DSM-5 dengan algoritma *Support vector machine* (SVM) dengan pendekatan CRISP-DM, dan penanganan ketidakseimbangan data dengan teknik *oversampling* SMOTE. Dataset terdiri dari data gejala anak dengan label ASD (ASD 1, ASD 2, ASD 3) dan Non-ASD. Hasil evaluasi menunjukkan bahwa model SVM tanpa SMOTE menghasilkan performa terbaik dengan akurasi keseluruhan 98,77%, serta *precision*, *recall*, dan *f1-scores* yang tinggi dan seimbang. Sementara, penerapan SMOTE meningkatkan sensitivitas pada kelas minoritas, namun menurunkan akurasi menjadi 90,12%. Oleh karena itu, model SVM tanpa SMOTE dinilai paling efektif. Model ini berpotensi sebagai alat bantu diagnosis awal, namun tetap memerlukan validasi klinis lebih lanjut.

Keywords: Autisme, DSM-5, *Support vector machine*, SMOTE, *Machine learning*, dan Klasifikasi.

I. INTRODUCTION

Autisme adalah gangguan pada perkembangan neuropsikiatri yang memengaruhi fungsi otak yang terlibat dalam kemampuan sosial, berkomunikasi, serta berperilaku [1]. Berdasarkan World Health

Organization (WHO) sekitar 1 dari 100 anak di seluruh dunia mengalami autisme [2]. Di Indonesia sendiri diperkirakan sekitar 2,4 juta anak yang hidup dengan ASD [3]. Penyebab autisme melibatkan berbagai faktor, baik genetik maupun lingkungan. Anak-anak yang didiagnosis autisme cenderung mengalami

kesulitan merespon orang lain, menghadapi kesulitan berkomunikasi dan berbahasa, serta menunjukkan reaksi yang tidak biasa terhadap lingkungan sekitarnya [4]. Proses diagnosis gangguan mental termasuk autisme mengacu pada panduan yang disebut *Diagnostic and Statistical Manual of Mental Disorders* edisi kelima (DSM-5), yang diterbitkan oleh *American Psychiatric Association* (APA). DSM-5 membagi tingkat keparahan ASD menjadi 3 level, yaitu Level 1 (membutuhkan dukungan), Level 2 (membutuhkan dukungan substansial), dan Level 3 (membutuhkan dukungan sangat substansial) [5]. Berdasarkan observasi dan wawancara yang dilakukan di Rumah Terapi dan Rumah Sakit terlihat bahwa mayoritas sebagian besar anak yang terdiagnosis ASD berada pada tingkat ASD Level 3 atau tingkat berat. Keterlambatan diagnosis pada anak autisme berdampak serius pada perkembangan anak, terutama dalam komunikasi, perilaku, dan sosial. Anak cenderung sulit mandiri dan butuh pendampingan jangka panjang, sehingga dapat menambah beban psikologis dan ekonomi keluarga.

Hal ini menunjukkan bahwa proses diagnosis dini masih lemah dan membutuhkan teknologi. Salah satu teknologi yang mendapat perhatian besar adalah machine learning. yang kini berkembang pesat di berbagai sektor, termasuk Kesehatan [6]. Algoritma machine learning mampu mengenali pola gejala autisme dari data anak dan memberikan potensi untuk diagnosis dini dan juga tepat [7]. Salah satu algoritma *machine learning* yang populer karena keakuratannya, terutama pada dataset kecil, adalah *Support vector machine* (SVM). Support vector machine (SVM) merupakan sistem pembelajaran yang memanfaatkan hipotesis dalam bentuk fungsi-fungsi linear pada fitur yang memiliki dimensi tinggi dan dilatih menggunakan algoritma pembelajaran berdasarkan prinsip teori optimasi [8].

Dalam proses klasifikasi menggunakan machine learning kualitas data memiliki pengaruh besar terhadap hasil yang diperoleh. Salah satu hal yang sering ditemui dalam pemrosesan data adalah ketidakseimbangan distribusi kelas. Dalam kasus ini jumlah anak dengan gejala ringan atau sedang jauh lebih sedikit dibandingkan dengan anak yang berada di tingkat berat. Ketidakseimbangan ini dapat mempengaruhi kinerja model dalam melakukan klasifikasi. Terdapat dua pendekatan yang dapat

dilakukan untuk mengatasi ketidakseimbangan kelas, yaitu pendekatan level data dan pendekatan algoritmik [9]. Ketidakseimbangan kelas akan ditangani menggunakan pendekatan level data dengan teknik *oversampling* SMOTE. *Synthetic Minority Over-Sampling Technique* (SMOTE) salah satu teknik oversampling yang digunakan untuk menangani ketidakseimbangan distribusi kelas dalam dataset dengan cara menghasilkan sampel sintetis dari kelas minoritas dengan K tetangga terdekatnya [10].

Penelitian mengenai klasifikasi autisme pernah dilakukan oleh Hazarika et al. (2022) yang menyimpulkan banyak penelitian yang belum memasukkan kriteria DSM-5, yang lebih komprehensif dibandingkan dengan DSM-4, dan dari beberapa algoritma algoritma SVM menjadi salah satu metode dengan performa terbaik, sehingga membuka peluang pengembangan model berbasis DSM-5 [11]. Penelitian selanjutnya Supriyadi dkk, (2022) yang menggunakan algoritma Naïve Bayes dan Support vector machine untuk memprediksi autisme menggunakan pendekatan CRISP-DM dengan hasil performa yang cukup tinggi [12]. Serta pada penelitian (Mulla et al., 2021) yang melakukan penelitian dengan mengkombinasi PCA dan SMOTE untuk meningkatkan performa pada beberapa algoritma, dimana algoritma KNN menghasilkan akurasi performa tertinggi [10].

Berdasarkan pentingnya diagnosis dini pada anak dalam penanganan autisme serta adanya perbedaan kesimpulan dari penelitian sebelumnya, penelitian ini dilakukan dengan mengacu pada kriteria DSM-5 menggunakan pendekatan CRISP-DM, algoritma *Support vector machine* (SVM) dan teknik *oversampling* SMOTE. Tujuan dari penelitian ini adalah untuk mengevaluasi performa SVM dalam mengklasifikasi autisme berdasarkan DSM-5 serta melihat dampak penggunaan SMOTE terhadap klasifikasi.

II. METHODS

2.1 Pengumpulan Data

Data yang digunakan dalam penelitian ini merupakan data primer, berupa rekam medis atau asesmen anak dengan ASD yang mencakup tingkat keparahan sesuai DSM-5, diperoleh dari dua institusi, yaitu Rumah Terapi dan Rumah Sakit. Dataset berisi gejala-gejala yang dialami

anak usia 2-15 tahun dalam kategori ASD Level 1, ASD Level 2, ASD Level 3, dan Non-ASD. Keseluruhan data telah dianonimkan dan dikodekan sesuai dengan prosedur etis. Selain itu, pengumpulan data didukung dengan observasi dan wawancara untuk memahami konteks dan validitas data yang digunakan.

2.2 Metode Implementasi Eksperimen

Penelitian ini menggunakan pendekatan CRISP-DM (*Cross-Industry Standard Process for Data Mining*), pada penelitian ini hanya sampai pada tahap evaluation, yaitu:

2.2.1 Business Understanding

Pada tahap ini dilakukannya pemahaman tujuan yang ingin dicapai pada penelitian. Penelitian ini bertujuan memanfaatkan machine learning untuk klasifikasi ASD sebagai alat bantu deteksi dini, mendukung intervensi lebih cepat. Model juga diharapkan memberi wawasan fitur signifikan dalam diagnosis. Hasil observasi dan studi pustaka digunakan untuk menentukan model yang sesuai agar akurasi klasifikasi meningkat dan bermanfaat bagi proses diagnosis.

2.2.2 Data Understanding

Penelitian ini menggunakan 81 data anak dari Rumah Terapi dan Rumah Sakit, yang mencakup anak dengan diagnosis ASD berdasarkan DSM-5 dan Non-ASD. Data diperoleh melalui observasi, rekam medis (asesmen) dan izin dari pihak terkait, terdiri dari 13 parameter yang mencerminkan aspek komunikasi, interaksi sosial, perilaku repetitif, dan faktor demografis sesuai DSM-5. Rincian parameter adalah sebagai berikut:

Tabel 1. Rincian Parameter

Parameter	Kategori & Skor
Kontak Mata	0 = Tidak, 1 = Baik
Kemampuan Berbicara	0 = Tidak Bisa, 1 = Terbatas, 2 = Lancar
Pemahaman Bahasa	0 = Tidak Memahami, 1 = Terbatas, 2 = Baik
Interaksi Sosial	0 = Tidak Ada, 1 = Terbatas, 2 = Baik
Stimming (Gerakan repetitif)	0 = Tidak, 1 = Ya
Rigid (Fleksibilitas)	0 = Fleksibel, 1 = Kaku

Ekolali (Mengulang kata)	0 = Tidak, 1 = Ya
Dreaming (Menunjukkan perilaku menyendiri)	0 = Tidak, 1 = Ya
Kesadaran Sosial	0 = Rendah, 1 = Sedang, 2 = Tinggi
Jenis Kelamin	L, P
Umur	2-15 Tahun
CARS (<i>Childhood Autism Rating Scale</i>)	0 = Normal (0-29), 1 = Rendah (30-36), 2 = Sedang (37-40), 3 = Tinggi (40-60)
Diagnosis	0 = ASD 1, 1 = ASD 2, 2 = ASD 3, 3 = Non-ASD

Penentuan kategori tingkat keparahan ASD didasarkan pada kriteria yang telah ditetapkan dalam DSM-5. Klasifikasi ASD dibagi menjadi 3 tingkat keparahan, yaitu Level 1 (Membutuhkan Dukungan), Level 2 (Membutuhkan Dukungan Substansial), serta Level 3 (Membutuhkan Dukungan Sangat Substansial).

2.2.3 Data Preparation

Tahap ini bertujuan untuk mempersiapkan data dari berbagai sumber, agar dapat digunakan secara optimal dalam proses klasifikasi autisme pada anak. Proses ini dilakukan di Google Colab dengan menggunakan pustaka *Python* seperti *pandas*, *numpy*, *matplotlib*, dan *seaborn* untuk memudahkan manipulasi, analisis, dan visualisasi data.

1. Data Selection

Data yang digunakan terdiri dari 81 data anak dengan 13 atribut, yaitu 11 data numerik dan 2 data kategorikal. Usia anak pada penelitian ini berkisar 2–15 tahun. Data disimpan dalam format CSV untuk keperluan analisis.

2. Data Transformation

Transformasi data dilakukan dengan menggunakan *LabelEncoder* untuk mengubah data kategorikal menjadi numerik, seperti pada kolom jenis_kelamin yang dikonversi dari “L” dan “P” menjadi 0 dan 1. Dan kolom diagnosis yang dikonversi dengan ASD 1 menjadi 0, ASD

2 menjadi 1, ASD 3 menjadi 2, serta NON-ASD menjadi 3.

3. Data Cleaning

Data cleaning bertujuan memastikan data bebas dari kesalahan atau nilai kosong. Dalam penelitian ini dilakukan pengecekan terhadap data hilang dan data duplikat. Hasilnya, pada dataset tidak ditemukan data kosong maupun duplikat, sehingga tidak diperlukan tindakan lanjutan.

Setelah itu, dilakukan pemisahan fitur (X) dan label (y). Kolom diagnosis dihapus dari fitur (X) dan digunakan sebagai label (y) untuk diprediksi oleh model.

4. Normalisasi Data

Normalisasi dilakukan menggunakan StandardScaler, yaitu metode yang mengubah setiap nilai fitur menjadi skala standar dengan mean = 0 dan standar deviasi = 1. Tujuannya agar seluruh fitur memiliki skala yang seragam dan membantu model bekerja lebih optimal.

2.2.4 Modelling

Pada tahap ini setelah dataset melalui proses persiapan, langkah selanjutnya yaitu melakukan proses pemodelan menggunakan dataset yang telah dipersiapkan sebelumnya. Penelitian ini menggunakan salah satu model machine learning, yaitu *Support vector machine* (SVM). Tahapan dalam proses ini mencakup pembangunan model SVM tanpa menggunakan SMOTE dan pembangunan model SVM dengan menggunakan SMOTE untuk menangani ketidakseimbangan data. Proses Modelling:

1. Split dataset

Dataset pada penelitian ini dibagi menjadi data pelatihan (*training set*) dan data pengujian (*test set*). Pembagian ini dilakukan dengan rasio 80:20, dimana 80% dari data akan digunakan sebagai data pelatihan dan 20% sisanya digunakan sebagai data pengujian [13]. Rasio ini digunakan sesuai dengan dataset yang berukuran kecil, serta memiliki keseimbangan ideal antara data untuk belajar dan data untuk menguji.

2. *Synthetic Minority Over-Sampling Technique* (SMOTE)
Synthetic Minority Over-Sampling Technique (SMOTE) adalah teknik

oversampling yang digunakan untuk mengatasi ketidakseimbangan kelas (*imbalanced dataset*) dalam *machine learning*, yang dapat memengaruhi kemampuan model dalam mengenali kelas minoritas [13]. Tidak seperti duplikasi data, SMOTE menghasilkan sampel sintesis baru untuk kelas minoritas, sehingga distribusi data menjadi lebih seimbang sebelum pelatihan. Secara spesifik, SMOTE memilih sebuah kelompok E_i dan tetangganya secara acak. Kemudian sampel baru dihasilkan menggunakan persamaan berikut, dimana E_j diambil dari himpunan tetangga terdekat [10]:

$$E_{baru} = E_i + (E_j - E_i) \delta$$

Dimana δ merupakan nilai yang dipilih secara acak dalam interval (0,1). Nilai δ berperan dalam menentukan jarak sampel sintesis yang dihasilkan berada dari titik awal E_i [10]. SMOTE hanya diterapkan pada data pelatihan, karena penerapannya pada data pengujian dapat menyebabkan kebocoran informasi yang membuat evaluasi model menjadi tidak valid [13]. Pada penelitian ini, SMOTE digunakan untuk menangani ketidakseimbangan pada kelas “diagnosis”, dengan membangun pipeline yang mencakup standarisasi data, oversampling menggunakan SMOTE, dan pelatihan model SVM dengan kernel linear.

3. Klasifikasi *Support vector machine*

Support vector machine (SVM) adalah metode yang kuat untuk membangun pengklasifikasi yang memiliki tujuan untuk membuat batas keputusan antara dua kelas yang memungkinkan label dari satu atau lebih vektor fitur. *Hyperplane* atau batas dari keputusan diorientasikan sejauh mungkin dari titik data terdekat dari masing-masing kelas, titik terdekat ini disebut *support vector* [14]. Identifikasi *hyperplane* yang dapat direproduksi adalah bagian dari proses pelatihan fungsi Keputusan SVM. Proses ini memaksimalkan jarak (*margin*) antara *support vector* dari kedua label kelas. Oleh karena itu, *hyperplane* optimal adalah yang memaksimalkan margin antar kelas. Setelah pembagian data, model SVM dengan kernel linear dilatih menggunakan data latih untuk mencari *hyperplane* terbaik yang memisahkan kelas dengan optimal [15]. *Hyperplane* optimal dapat didefinisikan sebagai berikut:

$$wx^T + b = 0$$

Dimana w sebagai *weight vector*, x sebagai vektor fitur *input*, dan b sebagai bias. w dan b akan memenuhi ketidaksetaraan berikut untuk semua elemen *training set*:

$$wx_i^T + b \geq +1 \text{ jika } y_i = 1$$

$$wx_i^T + b \leq -1 \text{ jika } y_i = -1$$

Tujuan *training* model SVM untuk menemukan w dan b sehingga *hyperplane* dapat memisahkan data dan memaksimalkan *margin* $\frac{1}{\|w\|^2}$ [14].

2.2.5 Evaluation

Pada tahap ini, dilakukannya evaluasi model yang bertujuan untuk dapat mengetahui seberapa baiknya kinerja model dengan menggunakan SVM dalam mengklasifikasi data. Dilakukan evaluasi menggunakan *K-Fold Cross-validation* dengan $k = 5$, *confusion matrix* dengan *cross-validation* yang telah diterapkan sebelumnya, serta AUC-ROC. Evaluasi ini dapat memberikan gambaran yang lebih menyeluruh mengenai kinerja model dalam klasifikasi data ASD.

1. Cross-validation

K-Fold Cross-validation adalah metode evaluasi yang populer dalam *machine learning* karena seimbang antara akurasi dan efisiensi komputasi. Teknik ini membagi dataset menjadi k bagian (*fold*) yang sama besar. Model dilatih sebanyak k kali, dengan setiap iterasi menggunakan satu *fold* sebagai data uji dan sisanya sebagai data latih. Hasil dari seluruh iterasi kemudian dirata-ratakan untuk memperoleh estimasi akurasi yang lebih stabil. Penelitian ini menggunakan nilai k adalah 5 karena memberikan keseimbangan antara bias dan varians, serta menghasilkan estimasi error yang lebih akurat [16].

2. Confusion Matrix

Confusion matrix adalah sebuah tabel berukuran 2×2 yang menyajikan perbandingan antara hasil prediksi model dengan data aktual yang tersedia. *Confusion matrix* digunakan sebagai alat evaluasi dalam credit scoring. Pada *confusion matrix* terdapat beberapa komponen sebagai berikut [17]:

- True Postive (TP), data sebenarnya positif dan diprediksi sebagai positif.
- False Positive (FP), data sebenarnya negatif tetapi diprediksi sebagai positif (kesalahan dalam mendeteksi negatif)

- False Negative (FN), data sebenarnya positif tetapi diprediksi sebagai negatif (kesalahan dalam mendeteksi positif)

- True Negative (TN), data sebenarnya negatif dan diprediksi sebagai negatif.

Matrik *confusion matrix* memungkinkan perhitungan metrik lain seperti *precision*, *recall*, dan *f1-scores*.

3. Accuracy

Accuracy atau akurasi menunjukkan seberapa besar presentase model menebak dengan tepat, baik berupa kasus positif maupun negatif. Berikut rumus perhitungan akurasi [18]:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

4. Precision

Precision atau presisi mengukur seberapa tepat model dalam memprediksi sesuatu sebagai positif. Berikut rumus perhitungan *precision* [18]:

$$Precision = \frac{TP}{TP + FP}$$

5. Recall

Recall dikenal sebagai sensitivitas yang fokus pada seberapa baik model mengklasifikasi seluruh kasus positif yang ada. Berikut rumus perhitungan *recall* [18]:

$$Recall = \frac{TP}{TP + FN}$$

6. F1-Scores

F1-Scores merupakan metrik yang menggabungkan *precision* dan *recall* ke dalam satu angka, menggunakan rata-rata harmonik. *F1-Scores* mempertimbangkan keseimbangan antara seberapa tepat model mengklasifikasi (*precision*) dan seberapa lengkap model menangkap semua kasus positif (*recall*). Berikut rumus perhitungan *f1-scores* [18]:

$$F1 - Scores = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

7. Receiver Operating Characteristic

Receiver Operating Characteristic (ROC) merupakan kurva atau grafik yang menggambarkan performa model klasifikasi. Pada grafik ini, sumbu x menunjukkan *false positive rate* (FPR) atau tingkat kesalahan positif, sedangkan sumbu y menunjukkan *true positive rate* (TPR) atau sensitivitas [13].

8. Area Under the Curve

Area Under the Curve (AUC) merupakan ukuran kuantitatif dari luas area di bawah kurva ROC. Nilai AUC diantara 0,5 sampai 1,0. Nilai 0,5 menunjukkan bahwa model memiliki performa seperti tebakan acak (random), sedangkan nilai 1,0 menunjukkan bahwa model mampu melakukan klasifikasi dengan sempurna [13].

III. RESULTS

3.1 Pengumpulan Data

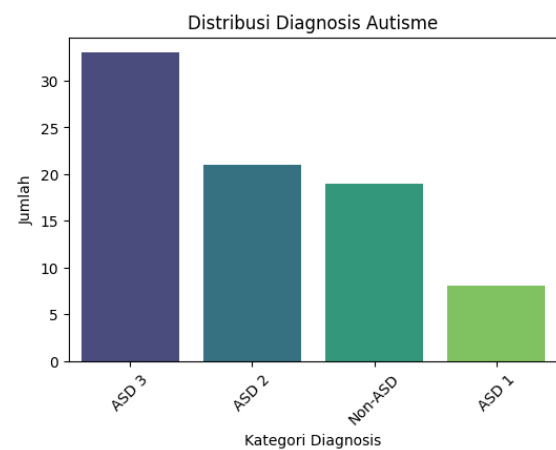
Pengumpulan data dilakukan dengan melibatkan observasi dan wawancara dalam mengambil data primer berupa rekam medis atau asesmen anak dengan ASD dan Non-ASD. Data yang digunakan sebanyak 81 data anak berisi gejala-gejala anak dengan diagnosis ASD sesuai kriteria DSM-5.

3.2 Analisis Distribusi Kelas Minoritas

Pada penelitian ini, label yang digunakan adalah “diagnosis” dengan jumlah data yang tidak seimbang. Berikut jumlah setiap kelas di sampel diagnosis:

Tabel 2. Jumlah Distribusi Diagnosis

Kategori	Jumlah
ASD 1	8
ASD 2	21
ASD 3	33
Non-ASD	19



Gambar 1. Distribusi Diagnosis Autisme

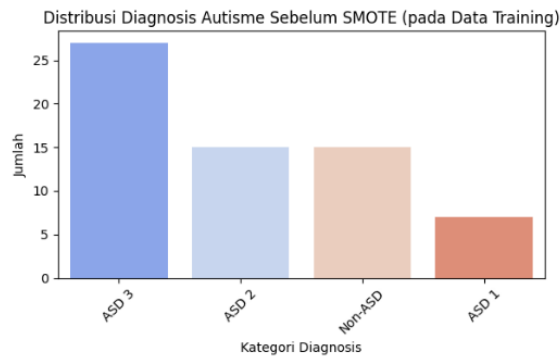
Pada distribusi diagnosis ASD mayoritas anak dalam dataset ini masuk ke dalam kategori ASD Level 3, dengan kategori ASD Level 1 memiliki jumlah data yang lebih sedikit dibandingkan dengan kategori lainnya. Ketidakseimbangan ini akan ditangani dengan menerapkan metode *oversampling* dengan menambah jumlah sampel pada kelas minoritas menggunakan *Synthetic Minority Over-Sampling Technique* (SMOTE) untuk memastikan model dapat mengklasifikasi ASD secara akurat tanpa bias pada kelas yang mayoritas.

3.3 Penanganan *Imbalanced Dataset*

Synthetic Minority Over-Sampling Technique (SMOTE) merupakan teknik *oversampling* yang digunakan pada machine learning untuk menangani ketidakseimbangan kelas (*imbalanced dataset*) pada data pelatihan. Berikut jumlah setiap kelas di sampel diagnosis pada data pelatihan:

Tabel 3. Jumlah Data Pelatihan

Kategori	Jumlah
ASD 1	7
ASD 2	15
ASD 3	27
Non-ASD	15

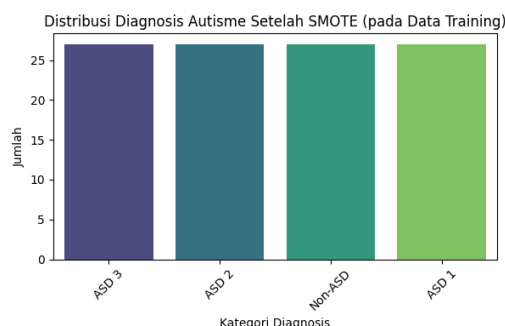


Gambar 2. Distribusi Diagnosis Data Pelatihan

Distribusi diagnosis dengan kondisi seperti ini dapat menyebabkan model menjadi bias terhadap kelas mayoritas, sehingga SMOTE perlu diterapkan untuk menambah jumlah sampel pada data yang lebih sedikit. SMOTE menggunakan $k_neighbors=1$ merupakan parameter yang menandakan bahwa SMOTE akan menggunakan satu tetangga terdekat untuk membuat data sintetis. Setelah proses pelatihan model dengan *pipeline* yang menggabungkan SVM dengan SMOTE, jumlah sampel pada setiap kelas diagnosis pada data pelatihan berubah menjadi lebih seimbang, yaitu:

Tabel 4. Jumlah Data Pelatihan Setelah SMOTE

Kategori	Jumlah
ASD 1	27
ASD 2	27
ASD 3	27
Non-ASD	27



Gambar 3. Distribusi Diagnosis Setelah SMOTE

pada Tabel 4 dan Gambar 3 terlihat bahwa distribusi diagnosis pada data pelatihan menjadi seimbang, sehingga model dapat belajar dengan lebih baik tanpa adanya bias terhadap kelas mayoritas.

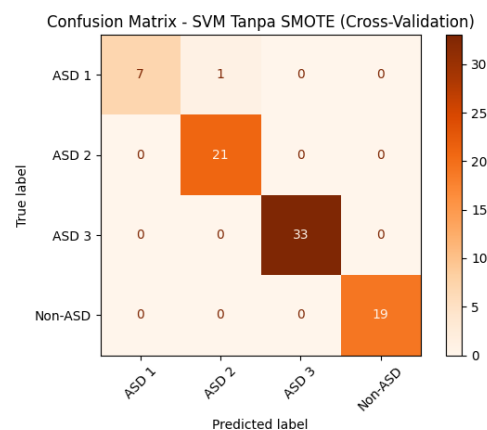
3.4 Hasil Evaluasi Model SVM Tanpa SMOTE

Evaluasi dilakukan terhadap model SVM tanpa menerapkan teknik *oversampling* SMOTE. Dilakukan evaluasi *5-fold cross-validation* untuk memaksimalkan data kecil dan menghindari *overfitting*.

Tabel 5. Hasil *Cross-validation 5-fold*

Fold	Akurasi
Fold 1	100%
Fold 2	100%
Fold 3	100%
Fold 4	93.75%
Fold 5	100%
Rata-rata (Mean)	98.75%

Cross-validation 5-fold memperoleh hasil dengan 3 *fold* mencapai akurasi 100%, dan 1 *fold* mendapat 93,75%, dengan rata-rata akurasi 98,75%. Hal ini menunjukkan model stabil, akurat, dan dapat melakukan generalisasi dengan baik. Evaluasi dilakukan dengan menggunakan *confusion matrix* yang didasarkan dengan hasil *cross-validation* sebelumnya.



Gambar 4. Hasil *Confusion matrix* SVM Tanpa SMOTE

Model menunjukkan tingkat *True Positive* (TP) yang tinggi pada kelas ASD 2, ASD 3, dan Non-ASD, menandakan kemampuan klasifikasi yang baik. Namun, pada kelas ASD 1 masih terdapat beberapa *False Positive* (FP) dan *False Negative* (FN), yang mengindikasikan kesalahan dalam membedakan kategori yang berdekatan. Meskipun demikian, secara keseluruhan model SVM menunjukkan performa yang baik dengan akurasi yang cukup tinggi. Berdasarkan hasil

confusion matrix sebelumnya dapat dilakukan perhitungan nilai *precision*, *recall*, *f1-scores* dan *accuracy*.

Tabel 6. *Classification Report*

Kelas	<i>Precision</i>	<i>Recall</i>	<i>F1-scores</i>	ROC-AUC
ASD 1	1.00	0.88	0.93	1.00
ASD 2	0.95	1.00	0.98	1.00
ASD 3	1.00	1.00	1.00	1.00
Non-ASD	1.00	1.00	1.00	1.00
<i>Accuracy</i>	0.9877			

Model SVM tanpa SMOTE mencapai akurasi sebesar 98,77%, sejalan dengan hasil *cross-validation* sebesar 98,75%. *Precision* tertinggi diperoleh pada kelas ASD 3 dan Non-ASD (100%), sedangkan *recall* secara umum tinggi, meskipun pada kelas ASD 1 sedikit lebih rendah (88%) karena sebagian data terklasifikasi ke kelas lain. *F1-scores* yang tinggi menunjukkan keseimbangan antara *precision* dan *recall*. Secara keseluruhan, model menunjukkan performa klasifikasi yang sangat baik meskipun terdapat sedikit penurunan pada kelas minoritas (ASD 1) akibat ketidakseimbangan data.

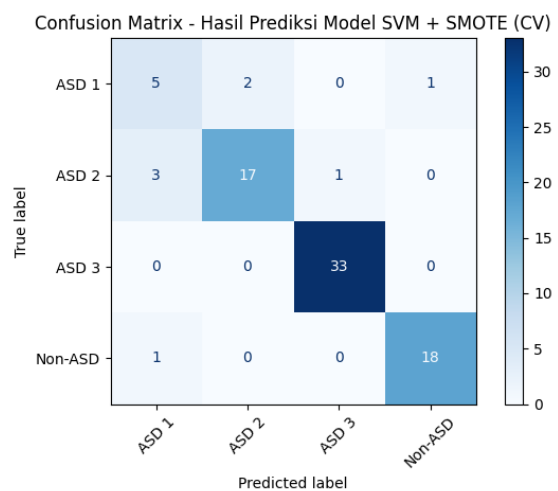
3.5 Hasil Evaluasi Model SVM Dengan SMOTE

Evaluasi dilakukan terhadap model SVM dengan menerapkan teknik *oversampling* SMOTE. Dilakukan *Cross-validation* 5-fold pada pipeline SMOTE + SVM.

Tabel 7. Hasil *Cross-validation* 5-fold

<i>Fold</i>	Akurasi
<i>Fold 1</i>	100%
<i>Fold 2</i>	100%
<i>Fold 3</i>	93.75%
<i>Fold 4</i>	87.50%
<i>Fold 5</i>	68.75%
Rata-rata (<i>Mean</i>)	90%

Tabel 7 menunjukkan hasil akurasi tinggi di sebagian besar *fold*, dengan *fold* 1 dan 2 mencapai 100%, namun *fold* ke-5 mengalami penurunan signifikan menjadi 68,75%, kemungkinan akibat dominasi sampel kelas minoritas atau variansi data. Rata-rata akurasi keseluruhan sebesar 90%. Evaluasi dilakukan dengan menggunakan *confusion matrix* yang didasarkan dengan hasil *cross-validation* sebelumnya.



Gambar 5. Hasil *Confusion matrix* SVM Dengan SMOTE

Model menunjukkan *True Positive* (TP) yang tinggi pada kelas ASD 2, ASD 3, dan Non-ASD, menandakan kemampuan klasifikasi yang baik. Namun, pada kelas ASD 1 masih terdapat beberapa *False Positive* (FP) dan *False Negative* (FN), terutama dari dan ke kelas ASD 2 dan Non-ASD, yang menunjukkan kesulitan dalam membedakan kategori yang berdekatan. Setelah menerapkan SMOTE, performa prediksi menjadi lebih merata, khususnya pada kelas minoritas seperti ASD 1 dan ASD 2, meskipun terjadi sedikit peningkatan kesalahan FP dan FN pada kelas tersebut. Secara keseluruhan, model SVM dengan SMOTE tetap menunjukkan kinerja klasifikasi yang baik dan akurasi yang tinggi. Berdasarkan hasil *confusion matrix* sebelumnya dapat dilakukan perhitungan nilai *precision*, *recall*, *f1-scores* dan *accuracy*.

Tabel 8. *Classification Report*

Kelas	<i>Precision</i>	<i>Recall</i>	<i>F1-scores</i>	AUC-ROC
ASD 1	0.56	0.62	0.59	0.98
ASD 2	0.89	0.81	0.85	0.99
ASD 3	0.97	1.00	0.99	1.00
Non-ASD	0.95	0.95	0.95	0.99
<i>Accuracy</i>	0.9012			

Model SVM dengan SMOTE mencapai akurasi 90,12%, konsisten dengan hasil *cross-validation* sebesar 90%. *Precision* dan *recall* tertinggi diperoleh pada kelas ASD 3 dan Non-ASD, masing-masing di atas 95%, menunjukkan kemampuan model mengenali kelas mayoritas dengan baik. Namun, pada kelas ASD 1, *precision* dan *recall* masih rendah,

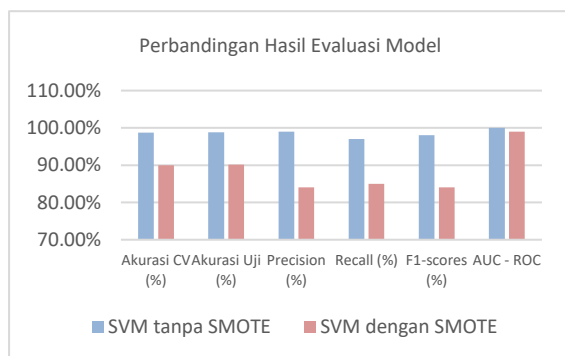
dan *F1-scores* hanya sebesar 59%, menandakan kesulitan model dalam membedakan kelas ini. Meskipun SMOTE membantu meningkatkan representasi kelas minoritas dan sensitivitas, performa keseluruhan model sedikit menurun dibandingkan SVM tanpa SMOTE. Namun, akurasi secara umum tetap tinggi.

3.6 Perbandingan Hasil Evaluasi

Setelah evaluasi model SVM tanpa dan dengan SMOTE menggunakan *cross-validation*, *confusion matrix*, *classification report*, serta ROC-AUC. Dilakukan perbandingan performa untuk menentukan model dengan akurasi dan kinerja keseluruhan yang lebih baik.

Tabel 9. Perbandingan Hasil Evaluasi Model

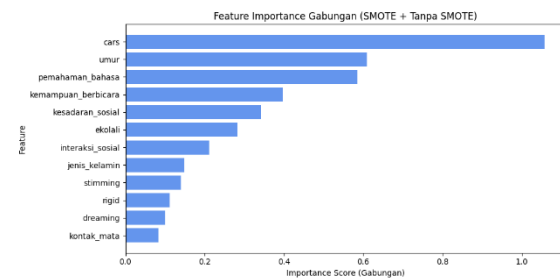
Metode	Akurasi CV (%)	Akurasi Uji (%)	Precision (%)	Recall (%)	F1-scores (%)	AUC - ROC
SVM tanpa SMOTE	98,75%	98,77%	99%	97%	98%	1.00
SVM dengan SMOTE	90%	90,12%	84%	85%	84%	0.99



Gambar 6. Perbandingan Hasil Evaluasi Model

Berdasarkan Tabel 9 dan Gambar 6, model SVM tanpa SMOTE menunjukkan akurasi *cross-validation* tertinggi sebesar 98,75%, sedangkan dengan SMOTE menurun menjadi 90,00%. Hasil ini konsisten dengan akurasi pada data uji, yakni 98,77% tanpa SMOTE dan 90,12% dengan SMOTE. *Precision*, *recall*, dan *F1-scores* juga lebih tinggi pada model tanpa SMOTE. Meskipun SMOTE menurunkan akurasi, namun dapat meningkatkan sensitivitas terhadap kelas minoritas seperti ASD 1. Secara keseluruhan, model SVM tanpa SMOTE dinilai sebagai model terbaik karena memberikan akurasi dan performa klasifikasi yang lebih seimbang.

3.7 Feature Importance



Gambar 7. Grafik *Feature Importance*

Berdasarkan metode permutation feature importance, pada Gambar 7 fitur paling berpengaruh dalam klasifikasi autisme adalah CARS dengan skor tertinggi 1.0572, menandakan perannya sangat penting dalam prediksi ASD. Fitur lain yang memiliki pengaruh signifikan adalah kemampuan_bicara, kesadaran_sosial, dan pemahaman_bahasa. Hasil ini menunjukkan bahwa model lebih banyak mempertimbangkan aspek klinis utama seperti CARS dan kemampuan komunikasi dibandingkan fitur perilaku minor dalam proses klasifikasi ASD.

IV. DISCUSSION

Hasil penelitian ini menunjukkan bahwa *support vvector machine* mampu mengklasifikasikan autisme sesuai dengan kriteria DSM-5 secara baik. Model SVM tanpa SMOTE menghasilkan kinerja yang lebih tinggi karena belajar menggunakan data asli dengan distribusi yang lebih alami. Sementara SMOTE membantu menyeimbangkan kelas, namun penerapan pada dataset kecil dapat menimbulkan *noise*. Dalam klinis ketepatan dalam klasifikasi sangat penting karena dapat mempengaruhi intervensi yang kurang tepat. Penelitian ini juga sejalan dengan studi sebelumnya yang menunjukkan bahwa meskipun SMOTE dapat meningkatkan sensitivitas terhadap kelas minoritas, akurasi model secara keseluruhan justru menurun [19]. Keterbatasan penelitian terletak pada ukuran dataset yang kecil, penerapan SMOTE yang menurunkan akurasi, dan hanya sampai tahap *evaluation*. Penelitian selanjutnya dapat menggunakan data yang lebih besar dan beragam dengan menggunakan teknik *sampling* lainnya, serta dapat mengembangkan dalam bentuk aplikasi mobile atau web.

V. CONCLUSION

Support vector machine (SVM) dalam klasifikasi autisme berdasarkan DSM-5 terbukti efektif dengan model SVM tanpa SMOTE yang menghasilkan akurasi tertinggi dan performa yang seimbang, menjadi pilihan terbaik dalam penelitian ini. Meskipun SMOTE mampu meningkatkan sensitivitas, namun hal ini menurunkan akurasi secara keseluruhan dengan selisih 8,65% dengan model SVM tanpa SMOTE yang dapat dikategorikan sebagai perbedaan yang cukup signifikan, terutama dalam konteks klasifikasi medis [20]. Penelitian ini menunjukkan potensi model SVM dalam membantu alat *skrining* awal autisme, namun implementasinya tetap harus melalui validasi klinis lebih lanjut.

BIBLIOGRAPHY

- [1] Dewi, S., & Morawati, S. (2024). Gangguan Autis pada Anak. *Scientific Journal*, 3(6), 418–431.
- [2] autism. (2023). Retrieved from World Health organization: <https://www.who.int/news-room/fact-sheets/detail/autism-spectrum-disorders>
- [3] Marini, S.KM, D. (2024, November 13). Kajian Epidemiologis, Anak dengan Autisme di Indonesia. Retrieved from Orang Tua Hebat: <https://www.orangtuahebat.id/kajian-epidemiologis-anak-autisme/>
- [4] Eko Cahyo, A., & Nilogiri, A. (2022). Klasifikasi Gangguan Autisme Pada Anak Menggunakan Algoritma C4.5 Dengan Teknik Random forest. *National Multidisciplinary Sciences*, 1(6), 846–851. <https://doi.org/10.32528/nms.v1i6.241>
- [5] APA. (2013). *Diagnostic And Statistical Manual of Mental Disorders Fifth Edition DSM-5*. In United States.
- [6] Zaenuddin, I., & Riyan, A. B. (2024). Perkembangan Kecerdasan Buatan (AI) Dan Dampaknya Pada Dunia Teknologi. *Jurnal Informatika Utama*, 2(2), 128–153.
- [7] Sugara, B., & Subekti, A. (2019). Penerapan Support vector machine (Svm) Pada Small Dataset Untuk Deteksi Dini Gangguan Autisme. *Jurnal Pilar Nusa Mandiri*, 15(2), 177–182. <https://doi.org/10.33480/pilar.v15i2.649>
- [8] Raj, S., & Masood, S. (2020). Analysis and Detection of Autism Spectrum Disorder Using Machine learning Techniques. *Procedia Computer Science*, 167(2019), 994–1004. <https://doi.org/10.1016/j.procs.2020.03.399>.
- [9] Widjiyati, N. (2021). Implementasi Algoritme Random forest Pada Klasifikasi Dataset Credit Approval. *Jurnal Janitra Informatika Dan Sistem Informasi*, 1(1), 1–7. <https://doi.org/10.25008/janitra.v1i1.118>
- [10] Mulla, G. A. A., Demir, Y., & Hassan, M. M. (2021). Combination of PCA with SMOTE Oversampling for Classification of High-Dimensional Imbalanced Data. *BEU Journal of Science*, 10(3), 858–869.
- [11] Hazarika, A., Kumar, C. J., & Das, P. R. (2022). Autism Spectrum Disorder diagnosis and machine learning: a review. *International Journal of Medical Engineering and Informatics*, 14(6), 512. <https://doi.org/10.1504/ijmei.2022.10050777>
- [12] Supriyadi, R., Maulidah, N., Fauzi, A., Nalatissifa, H., & Diantika, S. (2022). Penerapan Algoritma Naive Bayes Dan Support Vector Machine dalam Memprediksi Autisme. *Swabumi*, 10(1), 55–59. <https://doi.org/10.31294/swabumi.v10i1.12294>
- [13] Nigrou, N. (2024). Predicting Autism Spectrum Disorder Using Machine learning Classifiers. *January*. <https://doi.org/10.13140/RG.2.2.35833.44646>
- [14] Huang, S., Nianguang, C. A. I., Penzuti Pacheco, P., Narandes, S., Wang, Y., & Wayne, X. U. (2018). Applications of support vector machine (SVM) learning in cancer genomics. *Cancer Genomics and Proteomics*, 15(1), 41–51. <https://doi.org/10.21873/cgp.20063>
- [15] Pisner, D. A., & Schnyer, D. M. (2020). Support vector machine. In *Machine Learning: Methods and Applications to Brain Disorders*. Elsevier Inc. <https://doi.org/10.1016/B978-0-12-815739-8.00006-7>

- [16] Nti, I. K., Nyarko-Boateng, O., & Aning, J. (2021). Performance of Machine Learning Algorithms with Different K Values in K-fold Cross-Validation. *International Journal of Information Technology and Computer Science*, 13(6), 61–71. <https://doi.org/10.5815/ijitcs.2021.06.05>.
- [17] Zeng, G. (2020). On the confusion matrix in credit scoring and its analytical properties. *Communications in Statistics - Theory and Methods*, 49(9), 2080–2093. <https://doi.org/10.1080/03610926.2019.1568485>
- [18] Chicco, D., & Jurman, G. (2020). The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation. *BMC Genomics*, 21(1), 1–13. <https://doi.org/10.1186/s12864-019-6413-7>
- [19] Putri, V. M., Masjkur, M., & Suhaeni, C. (2021). Performance of SMOTE in a random forest and naive Bayes classifier for imbalanced Hepatitis-B vaccination status. *Journal of Physics: Conference Series*, 1863(1). <https://doi.org/10.1088/1742-6596/1863/1/012073>.
- [20] Sowjanya, A. M., & Mrudula, O. (2023). Effective treatment of imbalanced datasets in health care using modified SMOTE coupled with stacked deep learning algorithms. *Applied Nanoscience (Switzerland)*, 13(3), 1829–1840. <https://doi.org/10.1007/s13204-021-02063-4>

