

BAB II

KAJIAN LITERATUR

2.1 Artificial Intelligence (AI)

Artificial Intelligence (AI) merupakan bentuk kecerdasan buatan yang memanfaatkan kemampuan mesin atau komputer untuk meniru cara berpikir dan bertindak seperti manusia dalam menyelesaikan berbagai tugas yang mempermudah kehidupan. Perkembangannya terus menunjukkan peningkatan yang signifikan, seiring dengan pertumbuhan data yang sangat cepat serta dukungan algoritme yang semakin kompleks [9]. Selain itu, hadirnya berbagai kerangka kerja yang lebih terstruktur dan mudah diimplementasikan turut mempercepat proses pengembangan teknologi ini. Kemampuan sistem dalam mengolah dan memahami data secara lebih akurat, ditambah dengan kemajuan pesat pada aspek komputasi dan penyimpanan, membuat *Artificial Intelligence* (AI) terus berkembang dan tetap relevan hingga saat ini [10]. Kondisi ini menjadikan AI semakin banyak diterapkan di berbagai bidang, mulai dari industri hingga layanan publik.

Di sisi lain, *Artificial Intelligence* (AI) juga dapat dipahami sebagai sistem yang mengadaptasi kecerdasan manusia ke dalam perangkat komputer agar mampu meniru, menyimpan, serta mengelola informasi untuk membantu memenuhi kebutuhan sehari-hari. Teknologi ini dirancang agar komputer dapat menjalankan tugas yang membutuhkan proses berpikir, analisis, pengambilan keputusan, hingga pemecahan masalah sebagaimana yang biasa dilakukan manusia [11]. Dengan pendekatan tersebut, AI tidak hanya berfungsi sebagai alat bantu otomatisasi, tetapi juga sebagai pendukung dalam menghasilkan keputusan yang lebih cepat dan tepat. Penerapannya yang semakin luas menunjukkan bahwa AI memiliki peran penting dalam mendukung efisiensi dan produktivitas di berbagai sektor [10].

2.1.1 Machine Learning

Machine Learning merupakan salah satu bagian penting dalam kecerdasan buatan yang dimanfaatkan untuk menyelesaikan beragam permasalahan. Pendekatan ini dirancang agar mesin mampu mempelajari data dan meningkatkan kinerjanya secara bertahap tanpa harus diprogram secara rinci untuk setiap tugas. Dengan memanfaatkan algoritma tertentu, *Machine Learning* dapat mengenali pola serta hubungan yang kompleks di dalam data, tidak hanya bergantung pada aturan-aturan tetap yang telah ditentukan sebelumnya. Kemampuan tersebut memungkinkan pengambilan keputusan dilakukan dengan tingkat ketepatan yang

lebih baik [12]. Selain itu, teknologi ini juga mendorong efisiensi dalam pengolahan data berskala besar yang sulit dianalisis secara manual. Meskipun demikian, sistem *Machine Learning* memiliki kelemahan, terutama karena rentan terhadap penyusupan maupun manipulasi data melalui serangan yang dapat merugikan dan memengaruhi hasil analisis [13].

Secara umum, terdapat tiga paradigma utama dalam *Machine Learning* yang digunakan untuk mengekstraksi pola dan pengetahuan dari data, yaitu *Supervised Learning*, *Unsupervised Learning*, dan *Reinforcement Learning*. *Supervised Learning* merupakan metode pembelajaran yang menggunakan data berlabel sebagai acuan untuk melakukan prediksi atau klasifikasi. Berbeda dengan itu, *Unsupervised Learning* bekerja tanpa label sehingga sistem harus mengidentifikasi pola atau struktur data secara mandiri [14]. Sementara itu, *Reinforcement Learning* berfokus pada proses pembelajaran melalui mekanisme umpan balik berupa sinyal penguatan, yang membantu sistem menentukan tindakan paling optimal untuk mencapai tujuan tertentu [15]. Masing-masing paradigma memiliki karakteristik dan pendekatan yang berbeda, sehingga penerapannya perlu disesuaikan dengan jenis permasalahan dan hasil yang ingin dicapai. Perbedaan mendasar tersebut memberikan dampak yang signifikan terhadap implementasi di lapangan maupun kualitas keluaran yang dihasilkan [16].

2.1.2 Text Mining

Text mining merupakan salah satu cabang dari data *mining* yang berfokus pada proses ekstraksi informasi dan pengetahuan yang berguna dari data berbentuk teks tidak terstruktur. Berbeda dengan data numerik yang terorganisasi dalam basis data konvensional, data teks memerlukan tahapan praproses yang kompleks agar dapat diolah secara komputasional [15]. Proses *text mining* umumnya melibatkan beberapa tahap utama, seperti *text preprocessing* (terdiri dari tokenisasi, *stopword removal*, *stemming*, dan *lemmatization*), *feature extraction* (misalnya menggunakan metode *Bag of Words*, TF-IDF, atau *word embedding*), serta analisis lanjutan seperti *classification*, *clustering*, *sentiment analysis*, atau *topic modeling*. Melalui tahapan tersebut, sistem dapat mengenali pola tersembunyi dalam kumpulan dokumen teks dan mengubahnya menjadi representasi numerik yang dapat dipahami oleh algoritma pembelajaran mesin [17].

Secara umum, *text mining* banyak diterapkan dalam berbagai bidang seperti analisis opini publik di media sosial, pengelompokan dokumen, deteksi spam, *customer feedback analysis*, hingga sistem rekomendasi berbasis teks. Selain itu, perkembangan *Natural*

Language Processing (NLP) dan *Machine Learning* turut memperkuat kemampuan *text mining* dalam memahami konteks semantik dan sintaksis suatu bahasa. Dengan demikian, *text mining* menjadi teknologi penting dalam era big data untuk membantu pengambilan keputusan berbasis informasi tekstual secara lebih efisien dan akurat [17].

2.1.3 Natural Language Processing (NLP)

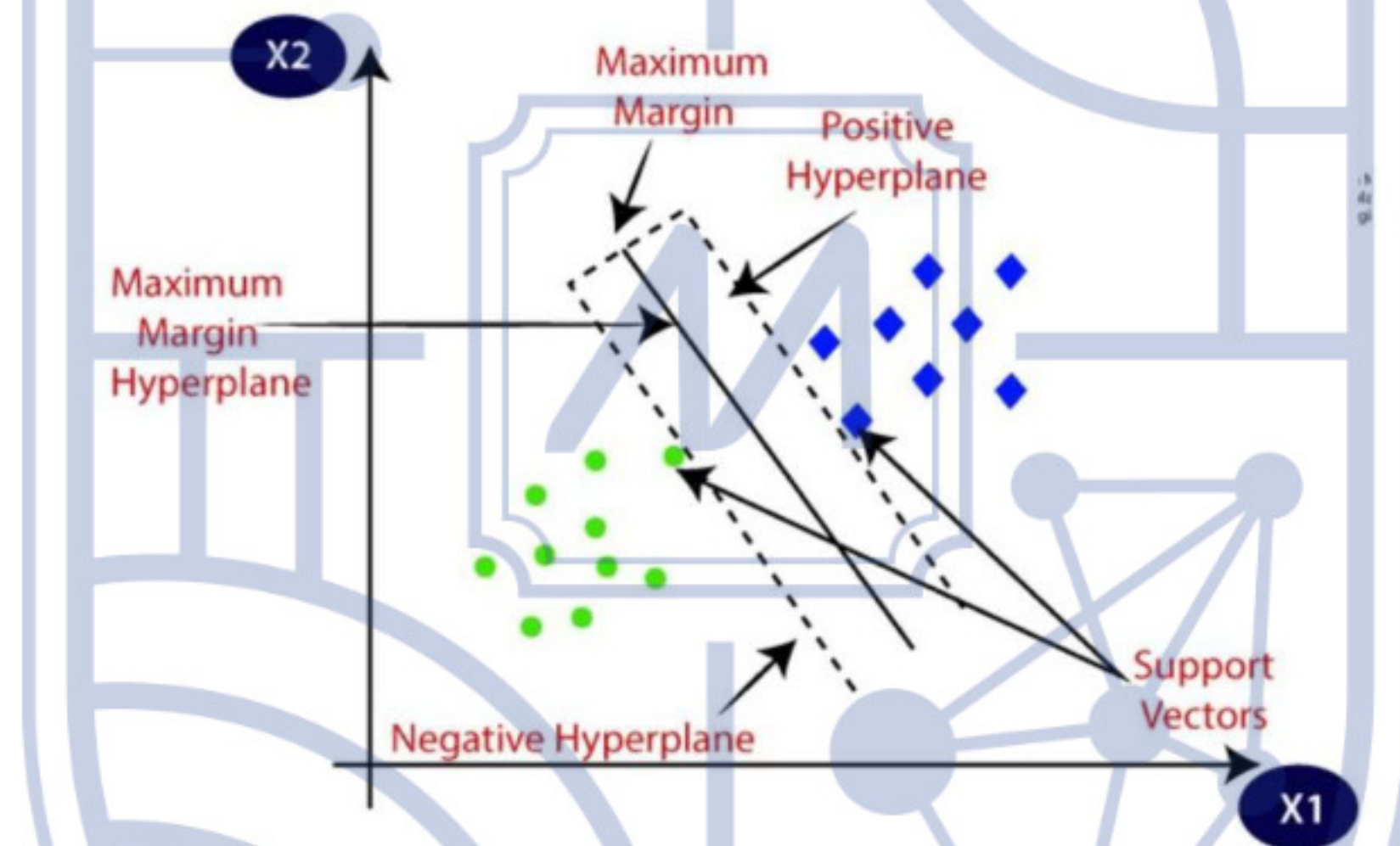
Natural Language Processing (NLP) merupakan cabang dari kecerdasan buatan (*Artificial Intelligence / AI*) yang berfokus pada kemampuan komputer untuk memahami, menafsirkan, dan memproses bahasa manusia secara alami. Tujuan utama NLP adalah menjembatani komunikasi antara manusia dan mesin dengan memungkinkan komputer mengenali makna dari teks atau ucapan sebagaimana manusia melakukannya [18]. Dalam praktiknya, NLP menggabungkan konsep dari *linguistics*, *computer science*, dan *machine learning* untuk menganalisis struktur bahasa, makna semantik, serta konteks komunikasi. Sederhananya, NLP memberikan akses dan sarana untuk analisis dari sebuah teks atau tulisan dengan membuat sebuah struktur yang baik dan teks yang sebelumnya belum terstruktur untuk dianalisis berikutnya. Kajian heuristik menyebutkan, NLP mulai diperkenalkan sekitar tahun 1970an dengan fungsi awal sebagai perangkat lunak dalam pemahaman bahasa domain block berlanjut pada dekade 80an dengan tambahan pengembangan dengan berbasis simbol dan ejaan [19].

2.2 Support Vector Machine (SVM)

Support Vector Machine (SVM) merupakan salah satu algoritma *machine learning* yang banyak digunakan untuk tugas klasifikasi dan regresi, termasuk dalam analisis teks. Konsep utamanya adalah menemukan garis atau bidang pemisah (*hyperplane*) yang mampu membedakan dua kelas data dengan margin paling besar sehingga proses klasifikasi menjadi lebih optimal. Semakin lebar jarak antar kelas yang terbentuk, semakin baik pula kemampuan model dalam melakukan generalisasi terhadap data baru [20]. Dalam bidang *text mining* dan *Natural Language Processing* (NLP), SVM kerap dimanfaatkan untuk tugas seperti deteksi spam, analisis sentimen, dan klasifikasi dokumen. Teks yang telah melalui tahapan preprocessing kemudian direpresentasikan dalam bentuk numerik, misalnya menggunakan TF-IDF atau *word embedding*, sebelum digunakan sebagai data latih pada model SVM [21].

Keunggulan SVM terletak pada kemampuannya menangani data berdimensi tinggi, dimana kondisi yang sangat umum dijumpai pada data teks. Algoritma ini berfokus pada

sejumlah titik data penting yang disebut *support vectors*, yaitu titik-titik yang paling berpengaruh dalam menentukan batas pemisah antar kelas. Selain itu, SVM dapat memanfaatkan berbagai fungsi kernel seperti *linear*, *polynomial*, dan *radial basis function* (RBF) untuk memetakan data yang tidak terpisah secara linear ke dalam ruang berdimensi lebih tinggi agar dapat dipisahkan dengan lebih baik. Karakteristik ini membuat SVM relatif tahan terhadap *overfitting* serta mampu memberikan tingkat akurasi yang baik pada data yang kompleks [22]. Meskipun metode *deep learning* semakin banyak digunakan dalam penelitian NLP, SVM tetap menjadi pilihan yang relevan karena efisiensinya dan kemudahan dalam menginterpretasikan hasil. Penggunaan *word embedding* yang tepat dikombinasikan dengan SVM dapat menghasilkan model klasifikasi teks yang cepat, stabil, dan memiliki performa yang kompetitif [23] [24].



Gambar 2.1 *Support Vector Machine* [24]

Menurut Burges proses klasifikasi dibedakan menjadi dua jenis berdasarkan kernel yang digunakan, yaitu SVM-Linear dan SVM Non-Linear [24].

1. SVM Linear

Pada proses klasifikasi menggunakan SVM Linear, data dapat dipisahkan dengan *hyperplane* menggunakan SVM (*Support Vector Machine*) dengan pembagian data yang jelas. Persamaan dari SVM kernel linear sebagai berikut: [25]

$$K(x_i, x_j) = x_i^T x_j \dots \dots \dots (1)$$

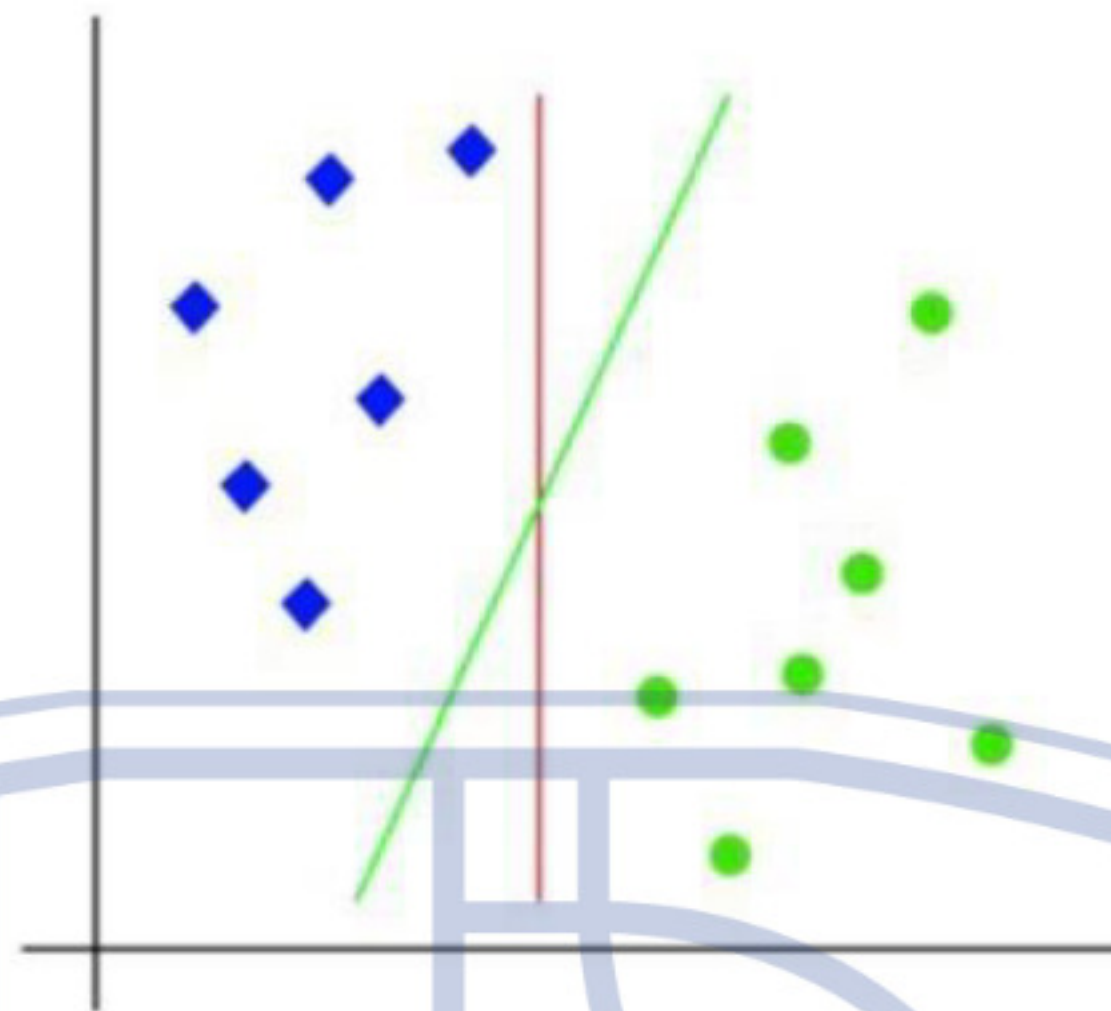
Keterangan:

$K(x_i, x_j)$: Fungsi kernel

x : Titik data masukan *Support Vector Machine*

i : Indeks atau nomor urutan data

T : Mengubah vektor baris menjadi vektor kolom dan sebaliknya



Gambar 2.2 SVM Kernel Linear [24]

2. SVM Non-Linear

SVM Non-Linear menggunakan fungsi kernel untuk memisahkan data non-linear menjadi data linear. Fungsi kernel juga dikenal sebagai trik kernel. Trik kernel digunakan untuk mengelompokkan data dari dimensi kecil menjadi dimensi besar. Ada beberapa fungsi kernel non-linear yang dapat digunakan dalam proses klasifikasi, antara lain:

a. Kernel polynomial

Kernel polynomial digunakan untuk merubah suatu data masukan ke dalam suatu dimensi yang memiliki ruang lebih tinggi. Kernel polinomial berfungsi untuk menemukan hyperplane yang baik untuk memisahkan data ruang yang telah diubah.

$$K(x_i, x_j) = (x_i^T x_j + C)^d \dots \dots \dots (2)$$

Keterangan:

- $K(x_i, x_j)$: Fungsi kernel
- x : Titik data masukan *Support Vector Machine*
- i : Indeks atau nomor urutan data
- T : Mengubah vektor baris menjadi vektor kolom dan sebaliknya
- C : Konstanta pada fungsi kernel
- d : pangkat atau derajat yang menentukan kompleksitas kernel

b. Kernel Gaussian RBF

Kernel RBF adalah kernel yang digunakan untuk mengklasifikasi data yang tidak dapat dipisah secara linear. Ruang fitur RBF memiliki jumlah dimensi yang tak terbatas yang ditentukan oleh parameternya. Oleh karena itu, ketika diulangi, dapat

menghasilkan solusi linear yang unik yang menjadikannya sebagai proses klasifikasi terbaik.

$$K(x_i, x_j) = \exp(-\gamma |x_i - x_j|^2), \gamma > 0 \dots \dots \dots (3)$$

Keterangan:

$K(x_i, x_j)$: Fungsi kernel

$\exp$: Fungsi eksponensial

γ : Parameter kernel yang mengontrol tingkat pengaruh jarak antar data

c. Kernel Sigmoid

Kernel sigmoid merupakan pengembangan dari JST (Jaringan Syaraf Tiruan). Kernel ini diusulkan secara teoritis untuk *Support Vector Machine* karena berasal dari jaringan syaraf. Kernel sigmoid umumnya bermasalah karena sulit untuk mendapatkan parameter dengan nilai positif [26].

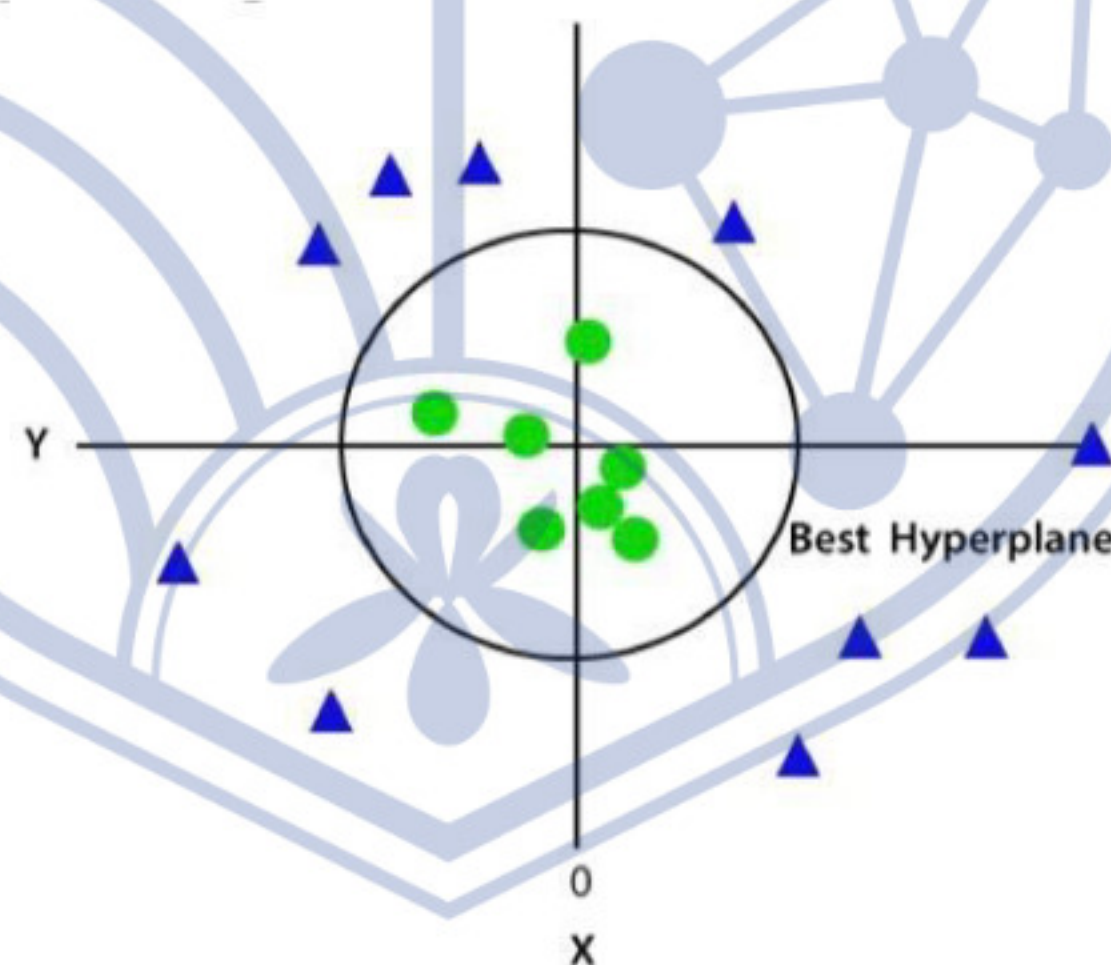
$$K(x_i, x) = \tanh \left(\gamma \left(x^T \frac{x}{\|x\|} \right) + r \right) \dots \dots \dots (4)$$

Keterangan:

$K(x_i, x)$: Fungsi kernel

$\tanh$: Fungsi aktivasi pada kernel

r : Konstanta pada fungsi kernel



Gambar 2.3 SVM Kernel Non-Linear [24]

Untuk memperoleh performa klasifikasi yang optimal, diperlukan proses optimasi hyperparameter pada model Support Vector Machine (SVM). Salah satu metode yang umum digunakan untuk melakukan optimasi tersebut adalah Grid Search. Grid Search merupakan metode pencarian hyperparameter yang bekerja dengan cara menguji seluruh kombinasi parameter yang telah ditentukan sebelumnya secara

sistematis untuk menemukan konfigurasi terbaik berdasarkan hasil evaluasi model. Pada algoritma Support Vector Machine (SVM), Grid Search dapat digunakan untuk mengevaluasi berbagai kombinasi hyperparameter seperti nilai C, gamma, dan jenis kernel. Setiap kombinasi parameter akan diuji dan dibandingkan menggunakan metrik evaluasi tertentu sehingga dapat diketahui konfigurasi parameter yang menghasilkan performa klasifikasi paling optimal. Melalui proses pencarian yang menyeluruh, Grid Search mampu membantu model memperoleh parameter yang lebih sesuai dengan karakteristik data yang digunakan[27].

Selain itu, penerapan Grid Search dapat meningkatkan kemampuan model dalam melakukan generalisasi terhadap data baru karena parameter yang digunakan telah melalui proses evaluasi dan pemilihan secara sistematis. Dengan demikian, penggunaan Grid Search tidak hanya membantu meningkatkan performa klasifikasi, tetapi juga mendukung pembentukan model SVM yang lebih optimal dan stabil untuk digunakan dalam proses deteksi komentar judi online pada platform YouTube[27].

Untuk melakukan klasifikasi, diperlukan pemodelan *Support Vector Machine* (SVM) pada data latih dan data uji. Data latih dihitung menggunakan metode analitik untuk menyelesaikan fungsi dual (*quadratic programming*). Berikut adalah Langkah Langkah yang harus dilakukan [28]:

- i. Melakukan penginputan data x , dan y
Inputan x dan y adalah nilai yang diambil dari hasil ekstraksi fitur indobERT yang digunakan untuk mencari hasil *feature x* dari model *Support Vector Machine* (SVM).

Keterangan:

x : Nilai vektor hasil ekstraksi

y : Nilai label

- ii. Inisialisasi parameter a

Selanjutnya, melakukan inisialisasi parameter a dengan kondisi awal $= 0$. Parameter a akan di optimisasi pada proses training untuk menentukan *Support Vector*.

$$a_i, a_j = 0$$

Keterangan:

a_i : Variable kelas ke - i

a_j : Variable kelas ke $-j$

iii. Menghitung kernel (*dot product*)

Selanjutnya, fungsi kernel dihitung dengan menggunakan dot product antara dua vektor data.

$$K(x_i, x_j) = x_i x_j \dots\dots\dots (5)$$

Keterangan:

x_i : Vektor data ke $-i$

x_j : Vektor data ke $-j$

iv. Membentuk fungsi dual Lagrangian (QP) untuk menentukan nilai parameter a

Selanjutnya, membentuk fungsi dual Lagrangian untuk mencari nilai parameter a .

$$L(a) = \sum_{i=1}^n a_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n a_i a_j y_i y_j K(x_i, x_j) \dots\dots\dots (6)$$

Keterangan:

$L(a)$: fungsi dual

n : jumlah data latih

y_i : Label data ke $-i$

y_j : Label data ke $-j$

$K(x_i x_j)$: Nilai kernel linear

Lakukan substitusi nilai Y dan K untuk melanjutkan perhitungan optimasi parameter a . Setelah melakukan substitusi. Maka akan ditetapkan constraint sebagai berikut.

$$\sum_{i=1}^n a_i y_i = 0$$

Setelah menetapkan constraint, akan dilakukan penyederhanaan fungsi dengan melakukan transformasi variabel bebas yang bergantung pada jumlah data yang akan dilatih untuk perhitungan. Kemudian, dilanjutkan dengan proses perhitungan optimasi parameter a .

v. Menghitung optimasi parameter a

Sehingga dapat ditentukan persamaan optimasi (turunan) parameter a adalah

$$\frac{\partial L}{\partial a_k} = 1 - \sum_{j=1}^n a_j y_k y_j K(x_k, x_j) = 0 \dots\dots\dots (7)$$

Keterangan:

$\frac{\partial L}{\partial a_k}$: turunan parsial fungsi lagrangian

vi. Menghitung nilai bobot (w)

Selanjutnya, menghitung bobot w untuk menentukan batas pemisah antar kelas.

$$w = \sum_{i=1}^n a_i y_i x_i \dots\dots\dots (8)$$

Keterangan:

a_i : hasil optimasi a ke - i

w : vektor bobot

vii. Menghitung nilai bias (b)

Nilai bias dihitung menggunakan salah satu vektor yang memenuhi kondisi $a > 0$. Maka nilai bias akan dihitung dengan data ke -1 pada *support vector*.

$$b = y_1 - w \times x_1 \dots\dots\dots (9)$$

Keterangan:

y_1 : label pada data ke -1

x_1 : vektor pada data ke -1

W : parameter *hyperplane* atau nilai bobot

viii. Model SVM (fungsi nilai Keputusan)

$$f(x) = w \cdot x_i + b \dots\dots\dots (10)$$

Keterangan:

w : parameter *hyperplane* atau nilai bobot

x_i : vektor data - i

b : nilai bias

Dengan aturan klasifikasi adalah sebagai berikut.

$f(x) > 0$ dinyatakan sebagai kelas +1 (judi)

$f(x) < 0$ dinyatakan sebagai kelas -1 (nonjudi).

Berdasarkan konsep dan mekanisme klasifikasi tersebut, pemilihan SVM dalam penelitian ini juga perlu mempertimbangkan karakteristiknya dibandingkan dengan metode klasifikasi lainnya. Perbandingan tersebut dilakukan untuk mengetahui kelebihan dan kekurangan SVM sehingga dapat menjadi dasar dalam menentukan kesesuaian metode SVM dengan kebutuhan penelitian. Berikut ini perbandingan SVM dengan metode klasifikasi lainnya sebagai berikut [29].

Tabel 2.1 Perbandingan SVM dengan Metode Lainnya

Metode	Kelebihan	Kekurangan
Support Vector Machine (SVM)	Akurasi tinggi pada data berdimensi tinggi, mampu bekerja baik pada dataset terbatas, tahan terhadap overfitting	Kurang efisien untuk dataset sangat besar, pemilihan kernel cukup kompleks,
Naive Bayes	Cepat, Sederhana, efisien dalam komputasi	Asumsi independensi fitur kurang realistik
Random Forest	Stabil, mampu menangani data dalam jumlah besar, tidak mudah overfitting	Model kurang mudah diinterpretasikan bisa menjadi kompleks

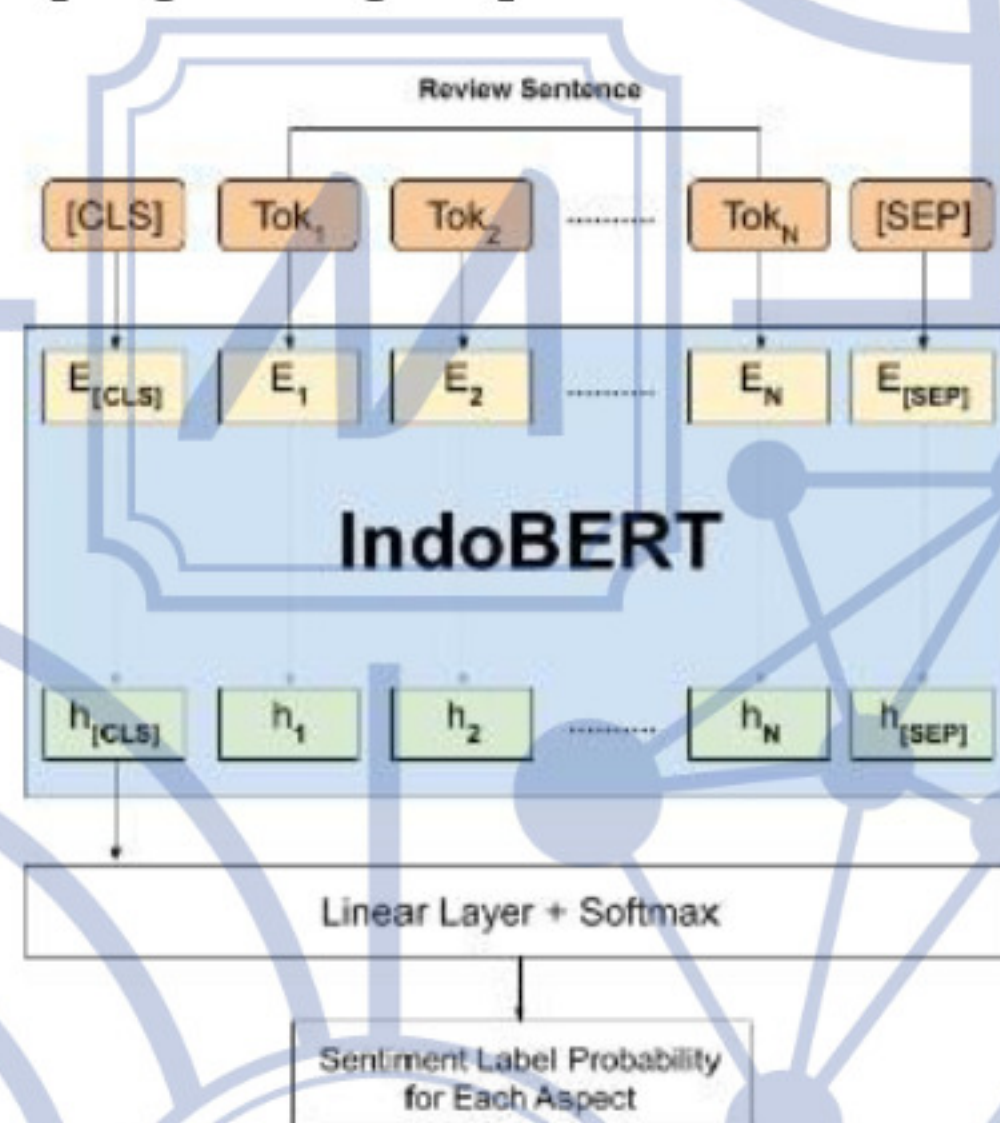
Alasan metode *Support Vector Machine* (SVM) dipilih karena memiliki kemampuan yang baik dalam mengklasifikasikan data dengan jumlah fitur yang kompleks meskipun jumlah data yang tersedia relatif terbatas. SVM juga dikenal memiliki tingkat generalisasi yang tinggi sehingga mampu mengurangi risiko overfitting dibandingkan metode lain yang membutuhkan data dalam jumlah besar. Selain itu, penggunaan kernel pada SVM memungkinkan pemodelan hubungan non-linear antar fitur secara lebih optimal, sehingga metode ini dinilai lebih efektif dan sesuai untuk kebutuhan penelitian dalam menghasilkan klasifikasi yang akurat dan stabil [29].

2.3 IndoBERT

IndoBERT merupakan model *language representation* berbasis arsitektur Transformer yang dirancang khusus untuk memproses dan memahami bahasa Indonesia. Model ini merupakan adaptasi dari BERT (*Bidirectional Encoder Representations from Transformers*) yang dikembangkan oleh Google, dengan perbedaan utama terletak pada data pelatihan yang sepenuhnya menggunakan korpus teks berbahasa Indonesia [30]. Tujuan utama pengembangan IndoBERT adalah agar sistem kecerdasan buatan dapat memahami konteks, struktur kalimat, dan makna kata dalam Bahasa Indonesia secara lebih akurat. Seperti BERT, IndoBERT menggunakan pendekatan *bidirectional context understanding*, yaitu menganalisis kata berdasarkan konteks di sebelah kiri dan kanan secara bersamaan,

sehingga pemahaman terhadap makna kalimat menjadi lebih menyeluruh dibandingkan model yang hanya membaca dari satu arah (*left-to-right* atau *right-to-left*) [31].

Proses pelatihan IndoBERT terdiri dari dua tahap utama, yaitu *Masked Language Modeling* (MLM), di mana sebagian kata dalam kalimat disembunyikan untuk diprediksi kembali oleh model, dan *Next Sentence Prediction* (NSP), yang melatih model untuk memahami hubungan antar kalimat [31]. Setelah melalui tahap pelatihan umum, IndoBERT dapat di-fine-tune untuk berbagai tugas spesifik seperti *text classification*, *sentiment analysis*, *named entity recognition*, dan *question answering*. Karena dilatih menggunakan data besar yang mencakup berbagai ragam teks berbahasa Indonesia seperti berita, artikel, dan media sosial, IndoBERT mampu menangani variasi bahasa formal maupun informal dengan baik. Dengan kemampuannya memahami konteks dan struktur bahasa yang kompleks, IndoBERT telah menjadi model dasar (*pre-trained* model) yang banyak digunakan dalam penelitian dan pengembangan aplikasi NLP berbahasa Indonesia [32] [33].



Gambar 2.4 Arsitektur Model IndoBERT [33]

Representasi teks atau *word embedding* merupakan metode untuk mengubah kata-kata menjadi vektor numerik agar dapat dipahami dan diolah oleh algoritma komputer. Pendekatan ini berbeda dengan representasi tradisional seperti *Bag of Words* atau TF-IDF yang hanya menghitung frekuensi kata tanpa memperhatikan makna atau konteksnya [34]. Dalam *word embedding*, setiap kata direpresentasikan sebagai vektor berdimensi tetap yang merefleksikan hubungan semantik antar kata, di mana kata-kata dengan makna serupa akan memiliki posisi vektor yang berdekatan di ruang vektor. Contohnya, kata “raja” dan “ratu” memiliki hubungan yang lebih dekat dibanding “raja” dan “mobil”. Melalui pelatihan pada korpus teks yang besar, model word embedding mampu menangkap pola linguistik dan hubungan semantik secara otomatis. Representasi ini kemudian digunakan sebagai masukan

(*input features*) untuk berbagai model *machine learning* atau *deep learning* pada tugas NLP seperti *text classification*, *sentiment analysis*, atau *named entity recognition* [35].

Selain itu, *word embedding* juga membantu mengurangi dimensi data teks yang sangat besar menjadi representasi yang lebih efisien tanpa kehilangan makna. Pada bahasa Indonesia, berbagai model *embedding* lokal telah dikembangkan agar mampu memahami struktur dan kosakata khas bahasa Indonesia. Dengan demikian, *word embedding* menjadi fondasi penting dalam analisis teks modern karena mampu menggabungkan makna semantik dan efisiensi komputasi dalam satu representasi yang komprehensif [35].

Dalam penelitian ini, IndoBERT digunakan untuk mengubah data komentar menjadi bentuk representasi yang dapat digunakan dalam proses klasifikasi. IndoBERT dipilih karena mampu memahami makna kata berdasarkan konteks kalimat, berbeda dengan Word2Vec dan GloVe yang menghasilkan representasi kata yang tetap tanpa mempertimbangkan konteks penggunaannya. Selain itu, IndoBERT telah dilatih menggunakan korpus berbahasa Indonesia sehingga lebih sesuai untuk mengolah komentar dalam bahasa Indonesia. Representasi dari keseluruhan kalimat diperoleh melalui token [CLS], kemudian digunakan sebagai fitur masukan untuk proses klasifikasi menggunakan Support Vector Machine (SVM)[36].

Dalam proses ekstraksi fitur, digunakan token khusus [CLS] sebagai representasi keseluruhan kalimat karena token ini dirancang untuk menangkap informasi global dari teks. Representasi [CLS] diperoleh setelah melalui proses *self-attention* sehingga mampu merangkum konteks secara menyeluruh dalam satu vektor. Penggunaan [CLS] lebih efisien dibandingkan metode lain seperti rata-rata seluruh token karena langsung menghasilkan representasi yang siap digunakan untuk klasifikasi [37]. Output dari IndoBERT yang digunakan adalah vektor berdimensi 768 yang merupakan standar pada model IndoBERT-base. Vektor ini kemudian digunakan sebagai input pada model SVM untuk menghasilkan klasifikasi yang lebih akurat dan stabil [38].

2.4 Confusion Matrix

Untuk mengukur tingkat akurasi suatu model klasifikasi, dapat digunakan teknik *confusion matrix*. *Confusion matrix* berukuran $N \times N$ berfungsi sebagai alat evaluasi kinerja model klasifikasi, di mana N merepresentasikan jumlah total kelas target. Matriks ini membandingkan antara nilai target aktual dan nilai target yang diprediksi oleh model *machine learning* [39]. Dalam *confusion matrix*, terdapat empat komponen utama yang

menggambarkan hasil proses klasifikasi secara menyeluruh. Keempat istilah tersebut yaitu [40]:

1. *True Positive* (TP) adalah jumlah kasus dimana sistem berhasil mendeteksi komentar yang mengandung unsur judi online dan mengklasifikasikan secara benar sebagai komentar berisi konten judi.
2. *True Negative* (TN) adalah jumlah kasus di mana sistem secara tepat mengenali komentar non-judi, yakni komentar yang tidak mengandung unsur judi online dan diklasifikasikan sebagai bukan komentar judi.
3. *False Positive* (FP) adalah jumlah kasus di mana sistem salah mengklasifikasikan komentar yang sebenarnya tidak mengandung unsur judi, namun terdeteksi sebagai komentar judi online.
4. *False Negative* (FN) adalah jumlah kasus di mana sistem gagal mendeteksi komentar yang sebenarnya mengandung unsur judi online, sehingga diklasifikasikan sebagai komentar non-judi [41].

		Actual values	
		+	-
Predicted values	+	True positive(TP)	False positive(FP)
	-	False negative(FN)	True negative (TN)

Gambar 2.5 *Confusion Matrix* [41]

Confusion matrix dapat digunakan untuk menghitung berbagai *performance metrics*. Hal ini bertujuan untuk mengukur kinerja model yang telah dibuat. Beberapa *performance metrics* yang umum digunakan yakni [33][40]:

1. *Accuracy*, menggambarkan seberapa akurat model dapat mengklasifikasikan dengan benar. Berikut ini adalah formula untuk menghitung *accuracy*.

$$\text{accuracy} = \frac{TP+TN}{TP+FP+TN+FN} \dots\dots\dots (11)$$

2. *Precision*, menggambarkan akurasi antara data yang diminta dengan hasil prediksi yang diberikan oleh model. Berikut ini formula untuk menghitung *precision*.

$$\text{precision} = \frac{TP}{TP+FP} \dots\dots\dots (12)$$

3. *Recall*, menggambarkan keberhasilan model dalam menemukan kembali sebuah informasi. Berikut ini formula untuk menghitung *recall*.

$$recall = \frac{TP}{TP+FN} \dots\dots\dots(13)$$

4. *F1 Score*, menggambarkan perbandingan rata-rata *precision* dan *recall* yang dibobotkan. Berikut ini formula untuk menghitung *F1 Score*.

$$F1\ score = \frac{2*recall*precision}{recall + precision} \dots\dots\dots(14)$$

Keterangan:

TP : *True Positive*

TN : *True Negative*

FP : *False Positive*

FN : *False Negative*

Jika *dataset* memiliki jumlah data *False Negative* dan *False Positive* yang sangat mendekati (*symmetric*), maka *accuracy* bisa digunakan untuk acuan performansi algoritme, sebaliknya *F1 Score* digunakan sebagai acuan apabila jumlah data *False Negative* dan *False Positive* tidak mendekati.

2.5 Media Sosial

Media sosial memiliki peranan penting bagi masyarakat, dengan melakukan berbagai kegiatan seperti pertukaran informasi maupun mencari informasi dengan lebih cepat, mudah, dan efektif. Saat ini, media sosial dapat dimanfaatkan untuk berbagai tujuan komunikasi yang banyak digunakan oleh organisasi, korporasi, pemerintah, hingga lembaga sosial masyarakat. Dengan kata lain, media sosial digunakan sebagai alat komunikasi dan telah menjadi kebutuhan primer bagi masyarakat, sehingga keberadaan media sosial semakin digemari oleh kalangan masyarakat baik anak-anak, dewasa, hingga kaum lanjut usia [42].

Media sosial merupakan media yang didesain untuk memudahkan interaksi sosial bersifat interaktif berbasis teknologi internet yang mengubah pola penyebaran informasi dari sebelumnya bersifat *broadcast media monologue* (satu ke banyak audiens) menjadi *social media dialogue* (banyak audiens ke banyak audiens). Media sosial telah menjadi bagian dari kehidupan manusia yang tidak terpisahkan. Hal ini dikarenakan media sosial mampu memenuhi kebutuhan manusia sebagai makhluk sosial dan selalu memiliki keinginan hadir sebagai sesuatu yang lain di mana hal ini tercermin dari pemilihan foto profil hingga pada cara manusia mencari kesenangan melalui media sosial [43].

YouTube merupakan salah satu platform berbagi video yang banyak digunakan oleh pengguna dari berbagai usia dan latar belakang. Platform ini memungkinkan pengguna untuk mengunggah, menonton, berdiskusi, serta membagikan video secara gratis. YouTube didirikan pada Februari 2005 oleh tiga mantan karyawan PayPal, yaitu Jawed Karim, Steve Chen, dan Chad Hurley. Sejak awal, YouTube menyediakan layanan untuk berbagai jenis video yang dapat diakses oleh pengguna. Melalui slogan “YouTube Broadcast Yourself”, platform ini pada awalnya bertujuan memberikan kesempatan kepada pengguna untuk membagikan dan menyimpan rekaman aktivitas mereka. Situs YouTube mulai dapat diakses melalui www.youtube.com pada 14 Februari 2005 dan kemudian mengalami perkembangan yang pesat hingga menjadi salah satu platform berbagi video yang banyak digunakan[44].

YouTube adalah sebuah situs daring yang menyediakan berbagai informasi dan menjadi wadah bagi semua orang untuk berbagi video secara *online* kepada orang lain. Situs ini khusus disediakan untuk mereka yang ingin mencari informasi dalam bentuk video dan menontonnya secara langsung. Pada zaman sekarang, YouTube semakin memiliki banyak fitur yang memudahkan penggunaannya untuk mencari informasi, seperti disediakan fitur pencarian sesuai dengan kategori seperti contohnya kategori musik, game, memasak, dan lain sebagainya. Selain itu, YouTube juga menyediakan fitur video siaran langsung yang di mana pemilik akun bisa berinteraksi dan berkomunikasi secara langsung dengan para penonton atau audiens, meskipun komunikasi tidak secara langsung tatap muka melainkan dengan perantara fitur chat [44].

2.6 Judi Online

Judi *online* dapat didefinisikan sebagai suatu bentuk kegiatan perjudian yang memanfaatkan jaringan internet sebagai sarana pelaksanaan taruhan dengan menggunakan uang sebagai alat pertaruhan. Fenomena ini semakin mudah berkembang karena kemajuan teknologi informasi yang memungkinkan akses tanpa batas waktu dan lokasi. Selain itu, lingkungan digital yang anonim memberi kesempatan bagi individu untuk mencoba taruhan tanpa merasa diawasi, sehingga risiko keterlibatan meningkat. Paparan terhadap iklan dan promosi daring, seperti bonus pendaftaran dan hadiah virtual, juga turut memikat pengguna baru, terutama remaja dan mahasiswa [45]. Di era digital saat ini, kejahatan siber dalam bentuk praktik perjudian *online* semakin marak terutama di kalangan mahasiswa dan remaja, karena berbagai alasan, termasuk rasa ingin tahu, tekanan teman sebaya, kebosanan, atau keinginan untuk menambah penghasilan. Kemudahan akses terhadap situs perjudian *online* berpotensi meningkatkan risiko individu untuk mengalami kecanduan judi [46].

Dalam dunia digital, penggunaan *platform* daring, media sosial, dan aplikasi mobile memfasilitasi penyebaran konten perjudian serta mempercepat interaksi pengguna dengan aktivitas taruhan, sehingga meningkatkan paparan terhadap konten perjudian dan risiko keterlibatan. Fenomena ini tidak hanya menimbulkan dampak sosial, ekonomi, dan psikologis, tetapi juga menghadirkan tantangan serius dalam upaya pencegahan dan penegakan hukum, mengingat sifat aktivitasnya yang tersembunyi dan lintas batas negara [47]. Selain itu, deteksi teks yang berkaitan dengan perjudian *online* turut menghadapi kendala teknis, seperti penggunaan istilah terselubung, konteks bahasa yang ambigu, serta perubahan pola komunikasi yang dinamis di media digital, sehingga menyulitkan sistem kecerdasan buatan untuk mengenali dan mengklasifikasikan konten perjudian secara akurat [48].

