

BAB II TINJAUAN LITERATUR

2.1 Analisis Sentimen

Analisis sentimen merupakan salah satu metode yang diterapkan dalam bidang *Natural Language Processing* (NLP) untuk mengidentifikasi, mengolah, serta mengklasifikasikan opini, perasaan, atau kecenderungan sikap yang terkandung dalam suatu data tekstual. (Bedir, 2021) Metode ini dapat dimanfaatkan untuk mengetahui kecenderungan opini masyarakat terhadap berbagai objek, seperti topik tertentu, produk, layanan, maupun informasi yang disampaikan melalui media digital. (Ligthart 2021) Menurut Ligthart, analisis sentimen merupakan kajian komputasional yang berfokus pada analisis opini, sentimen, emosi, dan sikap yang terkandung dalam teks yang dihasilkan oleh pengguna. Seiring dengan perkembangan teknologi informasi dan media digital, jumlah konten yang memuat opini dan tanggapan masyarakat semakin meningkat, sehingga diperlukan suatu metode yang mampu mengidentifikasi serta mengolah kecenderungan sentimen yang terdapat dalam data tersebut. Analisis sentimen memungkinkan data teks yang bersifat tidak terstruktur diolah menjadi informasi yang lebih sistematis sehingga dapat memberikan gambaran mengenai kecenderungan pendapat atau respons pengguna terhadap suatu objek tertentu.

analisis sentimen diterapkan untuk mengidentifikasi kecenderungan opini masyarakat terhadap informasi berita yang dipublikasikan oleh Kompas.com melalui platform Instagram. Data yang dianalisis berupa komentar masyarakat terhadap konten berita tersebut, yang selanjutnya diklasifikasikan ke dalam tiga kategori sentimen, yaitu positif, negatif, dan netral. Sentimen positif merujuk pada komentar

yang menunjukkan dukungan, persetujuan, apresiasi, atau tanggapan dengan kecenderungan yang bersifat positif. Sebaliknya, sentimen negatif menunjukkan adanya kritik, ketidaksetujuan, kekecewaan, atau tanggapan lain yang memiliki kecenderungan negatif. Sementara itu, sentimen netral merupakan komentar yang tidak memperlihatkan kecenderungan positif maupun negatif secara jelas. Hasil pengelompokan tersebut selanjutnya digunakan untuk menggambarkan kecenderungan respons masyarakat terhadap informasi berita yang menjadi objek penelitian. (Pratama 2022)

2.2 Pembelajaran Daring

Pembelajaran daring merupakan kegiatan pembelajaran yang memanfaatkan jaringan internet, LAN, dan WAN sebagai metode penyampaian, interaksi, serta fasilitas pembelajaran yang didukung oleh berbagai bentuk layanan belajar lainnya. (Anugrahana 2020) Pembelajaran daring merupakan proses pembelajaran yang dilaksanakan secara online dengan memanfaatkan teknologi internet, aplikasi pembelajaran, dan jejaring sosial sebagai sarana penyampaian materi, komunikasi, interaksi, serta pelaksanaan evaluasi pembelajaran. Pelaksanaan pembelajaran daring dilakukan tanpa tatap muka secara langsung antara pendidik dan peserta didik melalui berbagai platform pembelajaran yang tersedia. (Pratama et al. 2020)

Perkembangan teknologi informasi telah memberikan dampak yang besar dalam dunia pendidikan, terutama melalui penerapan pembelajaran daring. Perubahan ini semakin dipercepat oleh pandemi COVID-19 yang memaksa institusi pendidikan untuk beralih ke sistem pembelajaran digital. Namun, pelaksanaan pembelajaran daring menimbulkan berbagai pandangan dari pengguna, baik yang positif maupun negatif, sehingga diperlukan analisis yang sistematis sebagai evaluasi untuk perumusan kebijakan pendidikan Fajar Sidik et al., (2022).

Media sosial menjadi salah satu sumber utama dalam mengenali opini publik secara langsung. Platform seperti TikTok tidak hanya berfungsi sebagai sarana hiburan, tetapi juga telah bertransformasi menjadi media untuk menyebarkan konten edukatif. Komentar pengguna pada konten pembelajaran di TikTok mencerminkan pandangan, pengalaman, dan tingkat penerimaan terhadap pembelajaran daring, sehingga dapat dimanfaatkan sebagai data yang relevan untuk analisis menggunakan pendekatan *text mining* Apriani et al., (2024).

Analisis sentimen adalah salah satu teknik dalam Pemrosesan Bahasa Alami (NLP) yang digunakan untuk mengenali opini, sikap, atau emosi dalam teks. Metode ini memungkinkan klasifikasi opini ke dalam kategori tertentu, seperti positif, negatif, atau netral. Dalam konteks pembelajaran daring, analisis sentimen berfungsi untuk mengukur tingkat kepuasan dan mengidentifikasi masalah yang dihadapi oleh pengguna, sehingga dapat menjadi dasar bagi pengambilan keputusan yang didukung data Sarasvananda et al., (2022).

Salah satu algoritma yang banyak digunakan dalam analisis sentimen adalah Naive Bayes. Algoritma ini berbasis probabilitas dan memiliki keunggulan terkait kesederhanaan, efisiensi komputasi, serta kinerja yang baik dalam tugas klasifikasi teks. Prinsip kerja Naive Bayes melibatkan perhitungan probabilitas suatu dokumen termasuk ke dalam kategori tertentu berdasarkan distribusi kata-kata yang terkandung dalam dokumen tersebut Ernamia & Herliana, (2022).

Beberapa penelitian sebelumnya menunjukkan bahwa algoritma Naive Bayes memiliki tingkat akurasi yang kompetitif dalam analisis sentimen di media sosial. Misalnya, penelitian yang dilakukan oleh Friska Aditia Indriyani et al., (2023) menunjukkan bahwa Naive Bayes dapat menghasilkan klasifikasi sentimen yang efektif, bahkan dibandingkan dengan metode lain seperti *Support Vector Machine*

(*SVM*). Hal ini menandakan bahwa Naive Bayes adalah metode yang relevan untuk analisis data teks dalam volume besar.

Sejumlah penelitian menunjukkan bahwa penerapan analisis sentimen berbasis algoritma klasifikasi, seperti Naive Bayes, tidak hanya efektif dalam mengidentifikasi polaritas opini, tetapi juga mampu merepresentasikan tren persepsi pengguna secara agregat dalam konteks pembelajaran daring. Hal ini sejalan dengan temuan penelitian yang menyatakan bahwa data opini yang bersumber dari media sosial dapat diolah menjadi informasi strategis guna mendukung evaluasi kebijakan pendidikan, khususnya dalam memahami tingkat kepuasan serta berbagai kendala yang dialami pengguna selama berlangsungnya proses pembelajaran digital (Saputri et al., 2024).

Penelitian mengenai analisis sentimen dalam konteks pembelajaran daring umumnya menggunakan data dari platform seperti Twitter. Namun, pemanfaatan data dari TikTok masih cukup terbatas, padahal platform ini memiliki karakteristik unik berupa interaksi berbasis video pendek dan komentar yang lebih ekspresif. Oleh karena itu, penelitian ini mengusulkan pendekatan baru dengan menggunakan komentar di TikTok sebagai sumber data utama dalam analisis sentimen pembelajaran daring Dina et al., (2025).

Namun, komentar di media sosial biasanya bersifat singkat, tidak terstruktur, menggunakan bahasa sehari-hari, dan memiliki variasi makna berdasarkan konteks. Ciri-ciri ini menimbulkan kesulitan dalam proses analisis, sehingga dibutuhkan pendekatan komputasional yang dapat lebih mendalami konteks bahasa. Dalam penelitian ini, analisis sentimen dipakai sebagai metode untuk menilai kecenderungan pandangan masyarakat terhadap pembelajaran daring berdasarkan ungkapan emosional yang muncul dalam komentar Tiktok. Pratama et al., (2026)

Dengan demikian, penelitian ini bersifat tidak hanya analitis tetapi juga konstruktif karena mencakup perancangan sistem klasifikasi sentimen yang menggunakan algoritma Naive Bayes. Sistem yang dibuat diharapkan dapat menghasilkan model yang mampu secara otomatis mengevaluasi pandangan masyarakat terhadap pembelajaran daring. Selain itu, hasil penelitian ini diharapkan memberikan kontribusi praktis sebagai alat bantu dalam meningkatkan kualitas pembelajaran digital yang memanfaatkan media sosial. Apriani et al., (2024).

2.3 Tiktok

TikTok merupakan salah satu platform media sosial berbasis video pendek yang memungkinkan pengguna untuk membuat, mengunggah, membagikan, dan menikmati berbagai konten audiovisual. Platform ini menyediakan berbagai fitur pendukung, seperti musik, filter, serta fitur kreatif lainnya yang memberikan kesempatan kepada pengguna untuk mengekspresikan kreativitas dan menghasilkan konten secara menarik. Karakteristik TikTok yang berorientasi pada video pendek menjadikannya mudah digunakan dan memungkinkan penyampaian maupun penerimaan informasi secara cepat. Dalam perkembangannya, TikTok tidak hanya dimanfaatkan sebagai media hiburan, tetapi juga sebagai sarana komunikasi, ekspresi diri, kreativitas, dan penyebaran informasi kepada masyarakat. (Muhamad Nur Fitrianto 2025)

TikTok juga memiliki karakteristik interaktif yang memungkinkan pengguna memberikan respons terhadap konten melalui aktivitas menyukai, memberikan komentar, dan membagikan video. Platform ini memanfaatkan algoritma berbasis kecerdasan buatan untuk menganalisis preferensi pengguna berdasarkan aktivitas interaksi dan kemudian menyajikan konten yang dianggap relevan dengan minat pengguna. Dengan demikian, TikTok dapat dipandang sebagai media digital yang

tidak hanya berfungsi sebagai sarana konsumsi informasi, tetapi juga sebagai ruang interaksi dan pertukaran tanggapan antarpengguna. Kondisi tersebut menjadikan TikTok relevan sebagai media untuk mengamati berbagai respons dan kecenderungan opini pengguna terhadap informasi atau topik tertentu. (Wicaksono *et al.* 2024)

2.4 Text Mining

Text Mining adalah proses pemrosesan data teks untuk menemukan pola atau informasi yang berguna dengan menggunakan metode komputasi. Penambangan teks menjadi aspek penting dalam analisis sentimen karena data yang dianalisis biasanya berbentuk teks yang tidak terstruktur. Fokus utama dari *text mining* adalah mengidentifikasi pola, keterkaitan, atau pengetahuan baru yang tersembunyi dalam teks, yang kemudian dapat dimanfaatkan sebagai landasan untuk pengambilan keputusan.

Langkah-langkah umum dalam penambangan teks meliputi pengumpulan data, pra-pemrosesan, transformasi data menjadi bentuk numerik, pembuatan model, dan penafsiran hasil. Proses ini bertujuan untuk mengubah data teks menjadi informasi yang dapat dianalisis secara sistematis.

2.5 Preprocessing

Preprocessing teks merupakan tahapan awal dalam proses text mining yang bertujuan untuk membersihkan serta mempersiapkan data teks agar layak digunakan dalam tahap analisis selanjutnya. Tahapan ini mencakup beberapa proses utama, yaitu:

1. *Case folding*, yaitu proses mengubah seluruh huruf dalam teks menjadi huruf kecil untuk menyeragamkan format penulisan.

2. *Tokenization*, yaitu proses memecah teks menjadi unit-unit kata atau token. *Stopword removal*, yaitu proses menghilangkan kata-kata umum yang tidak memiliki kontribusi makna yang signifikan terhadap analisis.
3. *Stemming*, yaitu proses mengubah kata menjadi bentuk dasarnya.

Tahap preprocessing memiliki peranan yang sangat penting karena dapat meningkatkan kualitas data serta berkontribusi terhadap peningkatan akurasi model klasifikasi yang digunakan (Lestari & Hutagalung, 2025).

2.6 Natural Language Processing

Natural Language Processing (NLP) adalah bagian dari bidang kecerdasan buatan (*Artificial Intelligence*) yang menitik beratkan pada hubungan antara komputer dan bahasa manusia. NLP memungkinkan perangkat untuk menangkap, menginterpretasi, mengolah, dan menciptakan bahasa yang mendekati penggunaan manusia, baik dalam teks maupun suara (Sawicki, Ganzha & Paprzycki, 2023).

Dalam analisis sentimen, NLP memainkan peranan yang signifikan karena informasi yang diteliti berbentuk teks yang dibuat oleh pengguna. Proses NLP biasanya mencakup tahapan persiapan data seperti pembersihan, pengubahan huruf besar menjadi kecil, normalisasi kata, pemisahan kata, penghapusan kata umum, dan pengurangan kata ke bentuk dasar. Langkah-langkah tersebut bertujuan untuk meningkatkan mutu data sehingga lebih mudah ditangani oleh algoritme Machine learning maupun deep learning.

Dalam penelitian ini, NLP digunakan untuk menyiapkan data dari komentar Tiktok agar memiliki format yang konsisten sebelum dilakukan proses pelabelan sentimen dan klasifikasi dengan metode Algoritma Naive Bayes.

2.7 Machine Learning

Machine Learning merupakan salah satu cabang dari kecerdasan buatan (*Artificial Intelligence*) yang memungkinkan komputer mempelajari pola dari data untuk melakukan prediksi atau pengambilan keputusan tanpa harus diprogram secara eksplisit. Dalam machine learning, model dibangun berdasarkan data pelatihan (*training data*) sehingga mampu mengenali pola yang kemudian digunakan untuk mengklasifikasikan atau memprediksi data baru.

Pada bidang *Natural Language Processing* (NLP), machine learning banyak dimanfaatkan untuk berbagai tugas pengolahan teks, seperti klasifikasi dokumen, deteksi spam, analisis sentimen, dan pengenalan topik. Algoritma machine learning mampu mempelajari hubungan antarfitur pada data teks sehingga dapat mengelompokkan dokumen ke dalam kategori tertentu berdasarkan karakteristik kata-kata yang terkandung di dalamnya.

Dalam penelitian ini, pendekatan machine learning diterapkan untuk membangun model klasifikasi sentimen terhadap komentar pengguna TikTok mengenai pembelajaran daring. Model yang dikembangkan memanfaatkan algoritma Naive Bayes untuk mengelompokkan komentar ke dalam tiga kategori sentimen, yaitu positif, negatif, dan netral. Sebelum proses klasifikasi dilakukan, data teks terlebih dahulu melalui tahapan *preprocessing* sehingga menghasilkan representasi data yang lebih baik dan mampu meningkatkan kinerja model klasifikasi.

2.8 Term Frequency–Inverse Document Frequency (TF-IDF)

Term Frequency–Inverse Document Frequency (TF-IDF) merupakan metode pembobotan kata yang banyak digunakan dalam *text mining* dan *Natural Language Processing* (NLP). Metode ini bertujuan untuk mengubah data teks menjadi representasi numerik sehingga dapat diproses oleh algoritma *machine learning*. TF-IDF memberikan bobot yang lebih tinggi pada kata yang sering muncul dalam suatu dokumen tetapi jarang muncul pada keseluruhan kumpulan dokumen, sehingga kata tersebut dianggap lebih representatif dalam membedakan suatu dokumen dari dokumen lainnya.

TF-IDF terdiri atas dua komponen utama, yaitu *Term Frequency (TF)* dan *Inverse Document Frequency (IDF)*. Term Frequency digunakan untuk mengukur frekuensi kemunculan suatu kata dalam sebuah dokumen. Semakin sering suatu kata muncul pada dokumen tersebut, maka nilai TF akan semakin besar.

Rumus Term Frequency adalah sebagai berikut:

$$TF(t, d) = \frac{f(t, d)}{\sum_{t' \in d} f(t', d)}$$

Keterangan:

1. $TF(t, d)$ = nilai Term Frequency kata t pada dokumen d
2. $f(t, d)$ = jumlah kemunculan kata t pada dokumen d
3. $\sum_{t' \in d} f(t', d)$ = jumlah seluruh kata dalam dokumen d

Selanjutnya, *Inverse Document Frequency (IDF)* digunakan untuk mengurangi bobot kata-kata yang terlalu sering muncul pada seluruh dokumen karena kata tersebut dianggap kurang mampu membedakan karakteristik masing-masing dokumen.

Nilai IDF dihitung menggunakan persamaan berikut.

$$IDF(t) = \log\left(\frac{N}{df(t)}\right)$$

Keterangan:

1. N = jumlah seluruh dokumen
2. $df(t)$ = jumlah dokumen yang mengandung kata

Bobot akhir setiap kata diperoleh dengan mengalikan nilai TF dan IDF sehingga menghasilkan nilai TF-IDF sebagai berikut.

$$TF-IDF(t,d)=TF(t,d)\times IDF(t)$$

Melalui pembobotan tersebut, kata-kata yang memiliki kontribusi besar dalam membedakan suatu dokumen akan memperoleh nilai yang lebih tinggi dibandingkan kata-kata umum yang sering muncul pada seluruh dokumen.

Dalam penelitian ini, metode TF-IDF digunakan untuk mengubah komentar pengguna TikTok mengenai pembelajaran daring menjadi fitur numerik sebelum dilakukan proses klasifikasi menggunakan algoritma *Multinomial Naive Bayes*.

Representasi fitur menggunakan TF-IDF diharapkan mampu meningkatkan kualitas data masukan sehingga model dapat mengenali pola sentimen positif, negatif, maupun netral dengan lebih baik.(Zhang, 2025)

2.9 Algoritma Naive Bayes

Naive Bayes adalah algoritma klasifikasi yang didasarkan pada probabilitas, yang mengikuti Teorema Bayes dengan asumsi bahwa fitur-fitur bersifat independen satu sama lain(Systems et al., 2025). Algoritma ini sering digunakan dalam analisis sentimen karena kesederhanaannya dan kinerjanya yang cukup memuaskan.

Secara matematis, rumus Naive Bayes dapat dinyatakan sebagai berikut:

$$P(H | X) = \frac{P(X | H) \cdot P(H)}{P(X)}$$

Di mana $P(H|X)$ menunjukkan probabilitas hipotesis H berkaitan dengan data X , $P(X|H)$ adalah probabilitas data berdasarkan hipotesis, $P(H)$ adalah probabilitas awal dari hipotesis, dan $P(X)$ adalah probabilitas data.

Berdasarkan hal tersebut, Naïve Bayes dipilih sebagai metode utama dalam penelitian ini untuk mengelompokkan komentar Tiktok terkait pembelajaran daring ke dalam kategori positif, negatif, dan netral, sambil tetap mengutamakan efisiensi komputasi serta keandalan sistem

2.10 Penelitian Terdahulu

Untuk menguatkan dasar penelitian dan menentukan posisi penelitian yang dilakukan, penting untuk melakukan kajian terhadap penelitian-penelitian sebelumnya yang berkaitan dengan analisis sentimen di komentar Tiktok Tujuan dari kajian ini adalah untuk membandingkan objek yang diteliti, metode yang diterapkan, hasil yang didapatkan, serta keterbatasan dari masing-masing studi, sehingga bisa terdeteksi celah penelitian yang akan dilengkapi dalam penelitian ini. Rangkuman dari penelitian-penelitian sebelumnya yang relevan disajikan dalam Tabel 2.8

Table II .8. Penelitian Terdahulu

No	Penelitian	Tahun	Judul Penelitian	Metode	Data yang Digunakan	Hasil Penelitian
1	Apriani et al.	2023	Analisis Sentimen Penggunaan TikTok sebagai Media Pembelajaran	Naïve Bayes	Media sosial dalam pendidikan	TikTok dinilai efektif sebagai media pembelajaran dengan dominasi sentimen positif
2	Rahmadan i et al.	2022	TikTok Sentiment Analysis Using Naïve Bayes	Naïve Bayes	Data teks ulasan	Algoritma mampu mengklasifikasikan komentar dengan akurasi baik

3	Lestari & Hutagalung	2024	Evaluasi TF-IDF dalam Analisis Sentimen Naïve Bayes	SVM Naïve Bayes	Berbagai dataset teks	efektif pada data teks dengan preprocessing yang baik
4	Palomino & Aider	2022	Effectiveness of Text Preprocessing in Sentiment Analysis	Machine Learning	Data teks ulasan	Preprocessing meningkatkan akurasi model secara signifikan
5	Harahap et al.	2024	Sentiment Analysis Twitter Pemilu 2024	Naïve Bayes	Klasifikasi dataset	Model mampu mengolah data besar dan menghasilkan klasifikasi akurat

Berbagai penelitian terdahulu telah mengkaji penerapan analisis sentimen pada media sosial dengan menggunakan beragam objek penelitian dan metode klasifikasi. Penelitian (Apriani et al., 2023). menerapkan algoritma *Naïve Bayes Classifier* untuk menganalisis sentimen terhadap pemanfaatan TikTok sebagai media pembelajaran. Hasil penelitian menunjukkan dominasi sentimen positif, yang mengindikasikan bahwa TikTok mampu meningkatkan daya tarik dan interaktivitas pembelajaran. Namun, tingkat akurasi model masih dipengaruhi oleh kualitas data dan efektivitas tahapan *text preprocessing*.

(Rahmadani et al., 2022) juga menerapkan algoritma *Naïve Bayes* pada analisis sentimen komentar pengguna TikTok. Penelitian tersebut menunjukkan bahwa tahapan *text preprocessing*, seperti *case folding*, *tokenizing*, *stopword removal*, dan *stemming*, berperan penting dalam meningkatkan kualitas data sehingga menghasilkan klasifikasi sentimen yang lebih akurat. Temuan ini menegaskan bahwa keberhasilan klasifikasi dipengaruhi tidak hanya oleh algoritma, tetapi juga oleh kualitas pengolahan data.

(Lestari & Hutagalung, n.d. 2024)mengkaji penggunaan metode *Term Frequency–Inverse Document Frequency* (TF-IDF) sebagai teknik pembobotan kata dan membandingkan kinerja algoritma *Naïve Bayes* dengan *Support Vector Machine* (SVM). Hasil penelitian menunjukkan bahwa performa kedua algoritma dipengaruhi oleh karakteristik dataset dan kualitas *preprocessing*, sedangkan penggunaan TF-IDF mampu meningkatkan representasi fitur dan akurasi klasifikasi.

(Palomino, 2022) menjelaskan bahwa *text preprocessing* merupakan tahapan penting dalam analisis sentimen berbasis *machine learning*. Penerapan *case folding*, *tokenization*, *stopword removal*, *stemming*, dan normalisasi kata terbukti meningkatkan kualitas data sehingga menghasilkan performa klasifikasi yang lebih baik dibandingkan penggunaan data mentah.

Sementara itu, (Harahap et al., 2024)menunjukkan bahwa algoritma *Naïve Bayes* mampu mengolah data media sosial X (Twitter) dalam jumlah besar dengan tingkat akurasi yang baik. Hasil penelitian tersebut memperlihatkan bahwa algoritma *Naïve Bayes* masih efektif digunakan untuk menganalisis opini publik pada media sosial.

Berdasarkan berbagai penelitian terdahulu, dapat disimpulkan bahwa algoritma *Naïve Bayes* telah banyak digunakan dalam analisis sentimen dan menunjukkan performa klasifikasi yang baik. Namun, sebagian besar penelitian masih belum secara khusus mengkaji karakteristik komentar TikTok yang cenderung singkat, tidak baku, serta banyak mengandung singkatan, emoji, dan bahasa gaul. Karakteristik tersebut menjadi tantangan dalam proses *text preprocessing* dan klasifikasi sentimen. Oleh karena itu, penelitian ini berfokus pada analisis sentimen komentar pengguna TikTok menggunakan algoritma *Naïve Bayes* melalui tahapan *preprocessing* yang sistematis sehingga diharapkan dapat memberikan kontribusi terhadap pengembangan analisis sentimen pada data media social

2.11 Confusion Matrix

Confusion matrix adalah teknik evaluasi yang digunakan untuk menilai kinerja model klasifikasi dengan membandingkan hasil prediksi model dengan data sebenarnya. Confusion Matrix memberikan informasi tentang jumlah prediksi yang benar dan salah untuk masing-masing kategori yang diuji.

Dalam klasifikasi sentimen, confusion matrix digunakan untuk menilai seberapa baik model dapat membedakan antara sentimen positif, negatif, dan netral. Elemen utama dalam confusion matrix meliputi:

1. True Positive (TP), yaitu data yang diprediksi sebagai positif dan benar-benar positif.
2. True Negative (TN), yaitu data yang diprediksi sebagai negatif dan benar-benar negatif.

3. False Positive (FP), yaitu data yang diprediksi sebagai positif tetapi sebenarnya negatif.
4. False Negative (FN), yaitu data yang diprediksi sebagai negatif tetapi sebenarnya positif.

Melalui confusion matrix, dapat dihitung berbagai metrik evaluasi seperti *Accuracy*, *Precision*, *Recall*, dan *F1-Score* untuk menentukan sejauh mana keandalan model klasifikasi yang telah dibuat.

2.12 Accuracy

Accuracy adalah indikator yang digunakan untuk menilai seberapa tepat model dalam mengklasifikasikan seluruh data yang diuji. Nilai *accuracy* diperoleh dengan membandingkan jumlah prediksi yang benar dengan total data yang diuji.

Rumus *accuracy* adalah:

$$accuracy = \frac{TP + TN}{\text{Total data uji}}$$

Nilai akurasi yang lebih tinggi menunjukkan bahwa kemampuan model dalam mengklasifikasikan data semakin baik.

2.13 Precision

Precision adalah matrik yang digunakan untuk menilai seberapa akurat model dalam memprediksi kelas tertentu. *Precision* menggambarkan jumlah prediksi positif yang benar-benar sesuai dengan kenyataan.

Rumus *precision* adalah:

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

Nilai *precision* yang tinggi mengindikasikan bahwa model memiliki tingkat kesalahan False Positive yang rendah.

2.14 Recall

Recall adalah matrik yang digunakan untuk menentukan seberapa baik model dalam mendeteksi semua data yang benar dalam suatu kategori. *Recall* menunjukkan seberapa banyak data positif yang sukses dikenali oleh model.

Rumus recall adalah:

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

Nilai *recall* yang tinggi menandakan bahwa model mampu mengenali sebagian besar data yang seharusnya termasuk dalam kategori tertentu.

2.15 F1-Score

F1-Score adalah rata-rata harmonis dari *precision* dan *recall* yang digunakan untuk menggambarkan keseimbangan kinerja suatu model. Metrik ini sangat penting ketika terdapat ketidakseimbangan dalam distribusi data antar kelas.

Rumus F1-Score adalah:

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

Nilai F1-Score yang tinggi menandakan bahwa model memiliki keseimbangan yang baik antara *precision* dan *recall* dalam proses klasifikasinya.