



Skripsi

Analisis Sentimen pada Ulasan Aplikasi Stockbit Menggunakan Algoritma Naïve Bayes

**Jenis Skripsi:
Penelitian Eksperimental**

Disusun Sebagai Salah Satu Syarat Memperoleh Gelar Sarjana Komputer (S.Kom.)

Disusun oleh:

Albertus Andhika Dewa Satria Bintang Budaya
NIM. 20.0504.0037

Pembimbing:

Nuryanto, S.T., M.Kom.
NIDN. 0605037002

Pembimbing:

Maimunah, S.Si., M.Kom.
NIDN. 0612117702

**Program Studi Teknik Informatika
Fakultas Teknik
Universitas Muhammadiyah Magelang
Tahun 2025**

Bab 1 Pendahuluan

1.1 Latar Belakang Masalah

Saham merupakan salah satu instrumen pasar keuangan yang paling diminati oleh investor ritel maupun institusional karena potensinya dalam memberikan imbal hasil yang tinggi. Disisi lain, bagi perusahaan, saham menjadi salah satu alternatif utama dalam memperoleh pendanaan melalui pasar modal. Seiring berkembangnya teknologi digital, ekosistem investasi saham pun mengalami transformasi signifikan. Salah satu dampaknya adalah munculnya aplikasi finansial berbasis digital yang menyederhanakan proses transaksi dan analisis saham seperti Stockbit. Menurut data dari Otoritas Jasa Keuangan (OJK), jumlah investor ritel pasar modal di Indonesia mengalami pertumbuhan lebih dari 50% dalam dua tahun terakhir, seiring meningkatnya akses publik terhadap aplikasi investasi yang tidak hanya fungsional, tetapi mudah dipercaya. Di tengah kondisi ini, ulasan pengguna di Google Play Store menjadi sumber data penting yang mencerminkan kepuasan, kritik, maupun harapan pengguna terhadap aplikasi seperti Stockbit (Otoritas Jasa Keuangan, 2022).

Salah satu metode yang banyak digunakan dalam analisis sentimen adalah algoritma Naïve Bayes, karena memiliki keunggulan dalam hal kecepatan, kesederhanaan, dan efektivitas dalam mengolah data teks. Naïve Bayes bekerja berdasarkan prinsip probabilitas dengan mengasumsikan independensi antar fitur, yang membuatnya cukup andal dalam proses klasifikasi teks, termasuk ulasan aplikasi. Selain itu, beberapa penelitian terdahulu juga telah membuktikan keandalan Naïve Bayes dalam domain serupa.

Menurut studi yang dilakukan oleh (Rahel Lina Simanjuntak, Theresia Romauli Siagian, Vina Anggriani, & Arnita Arnita, 2023), algoritma Naïve Bayes menunjukkan performa yang baik dalam mengklasifikasikan sentimen pada ulasan aplikasi *e-commerce* dengan tingkat akurasi mencapai 85% hal ini menunjukkan potensi besar metode tersebut untuk diterapkan pada ulasan aplikasi finansial seperti Stockbit.

Analisis sentimen merupakan bagian dari cabang keilmuan Natural Language Processing (NLP) yang berfokus pada pemahaman opini, emosi, dan persepsi yang terkandung dalam suatu teks. Teknik ini digunakan untuk mengidentifikasi dan mengklasifikasikan sentimen menjadi kategori seperti positif, negatif, atau netral, dan sangat efektif dalam mengevaluasi opini publik dari data tidak terstruktur seperti ulasan aplikasi. Dalam penelitian yang dilakukan oleh (Jazuli, 2024) dijelaskan bahwa penerapan analisis sentimen yang dioptimalkan dengan model BERT mampu memberikan pemetaan yang lebih komprehensif terhadap umpan balik pengguna dalam bahasa Indonesia, termasuk dalam konteks layanan digital.

Studi serupa juga dilakukan oleh (Serlina & Rahim, 2025) yang menganalisis klasifikasi sentimen teks ulasan pengguna aplikasi Duolingo dalam bahasa Indonesia. Hasil penelitian mereka menegaskan bahwa pendekatan NLP seperti Naïve Bayes maupun BERT dapat mengungkap kecenderungan perilaku dan preferensi pengguna secara lebih akurat dibandingkan analisis manual, terutama dalam menghadapi data besar yang bervariasi secara emosional. Hal ini membuktikan bahwa analisis sentimen tidak hanya bermanfaat dalam sektor *e-commerce* atau sosial media, tetapi juga memiliki potensi besar dalam memahami kepuasan pengguna terhadap

aplikasi keuangan seperti Stockbit.

Aplikasi Stockbit ini menawarkan beberapa fitur, seperti analisis saham, diskusi komunitas, serta layanan transaksi yang memudahkan investor dalam mengambil keputusan investasi. Namun dengan meningkatnya jumlah pengguna, berbagai ulasan dan opini ini juga semakin banyak bermunculan, terutama di platform digital, seperti Google Play Store dan App Store. Stockbit menjadi salah satu aplikasi investasi dengan tingkat unduhan dan rating yang tinggi di Indonesia. Berdasarkan data Google Play Store, Stockbit telah diunduh lebih dari 1 juta kali dengan rating 4,7 dari 5 (Google Play Store, 2025). Namun, dibalik angka tersebut, terdapat variasi opini pengguna yang cukup signifikan. Beberapa pengguna memberikan ulasan positif mengenai fitur komunitas dan edukasi investasi, namun tidak sedikit pula yang mengeluhkan masalah teknis seperti error saat transaksi, keterlambatan update harga saham, serta tampilan aplikasi yang membingungkan bagi pemula. Masalah-masalah tersebut jika tidak ditangani dengan baik, dapat berdampak pada penurunan kepercayaan investor serta menjadi hambatan dalam akuisisi pengguna baru. Dalam situasi ini, analisis sentimen menjadi sangat penting untuk memahami secara menyeluruh bagaimana persepsi pengguna terhadap aplikasi. Berbeda dengan aplikasi lain seperti Bibit atau Ajaib yang fokus pada reksa dana dan robo advisor, Stockbit memiliki keunikan tersendiri karena memadukan fitur sosial media investor (komunitas diskusi), analisis teknikal, dan transaksi saham dalam satu aplikasi. Hal ini menjadikan pengalaman pengguna stockbit lebih kompleks, sehingga ulasan yang muncul pun cenderung lebih beragam dan dinamis.

Selain itu, berdasarkan observasi awal terhadap distribusi data ulasan pengguna, ditemukan adanya ketidakseimbangan jumlah antara ulasan positif dan negatif. Ketimpangan ini dapat menyebabkan model klasifikasi cenderung bias terhadap kelas mayoritas, sehingga akurasi terhadap kelas minoritas menjadi rendah. Permasalahan ini merupakan tantangan umum dalam analisis sentimen, terutama ketika jumlah data negatif jauh lebih sedikit dibanding data positif. Untuk mengatasi hal ini, berbagai teknik resampling telah digunakan, di antaranya SMOTE (*Synthetic Minority Over-sampling Technique*) dan ADASYN (*Adaptive Synthetic Sampling*). Kedua metode ini telah terbukti efektif meningkatkan performa model klasifikasi dalam kondisi data tidak seimbang, termasuk dalam studi analisis sentimen. Oleh karena itu, penelitian ini mengadopsi pendekatan tersebut untuk memastikan evaluasi sentimen yang lebih adil dan akurat terhadap ulasan aplikasi Stockbit. (Pamuji & Putri, 2023)

1.2 Perumusan Masalah

Berdasarkan latar belakang permasalahan di atas, maka masalah penelitian ini dirumuskan sebagai berikut :

- 1) Bagaimana proses klasifikasi sentimen ulasan pengguna aplikasi Stockbit dilakukan dengan menggunakan algoritma Naïve Bayes dioptimalkan dengan teknik resampling SMOTE dan ADASYN?
- 2) Bagaimana pengaruh penerapan teknik resampling SMOTE dan ADASYN terhadap akurasi algoritma Naïve Bayes dalam mengklasifikasikan sentimen ulasan pengguna aplikasi Stockbit?

- 3) Bagaimana perbandingan kinerja model Naïve Bayes dengan dan tanpa penerapan teknik resampling dalam penerapan teknik resampling dalam mengklasifikasikan sentimen ulasan pengguna aplikasi Stockbit?

1.3 Tujuan

“Berdasarkan rumusan masalah, maka tujuan penelitian yang akan dicapai adalah :

1. Menerapkan algoritma Naïve Bayes dalam melakukan klasifikasi sentimen terhadap ulasan pengguna aplikasi Stockbit, baik tanpa maupun dengan penerapan teknik resampling (SMOTE dan ADASYN) guna menangani ketidakseimbangan data.
2. Menganalisis pengaruh penerapan teknik resampling SMOTE dan ADASYN terhadap kinerja model klasifikasi Naïve Bayes, khususnya dalam meningkatkan akurasi, presisi, *recall* dan F1-score.
3. Menganalisis dan membandingkan efektivitas kinerja model Naïve Bayes dalam klasifikasi sentimen ulasan pengguna aplikasi Stockbit antara penerapan teknik resampling dan tanpa penerapan teknik resampling..

1.4 Manfaat

Berdasarkan tujuan penelitian yang telah disebutkan di atas, maka hasil penelitian ini dapat memberikan manfaat sebagai berikut:

1. Manfaat teoritis

Secara teoritis, penelitian ini memberikan kontribusi dalam pengembangan ilmu pengetahuan, khususnya di bidang data mining, text mining, dan analisis sentimen. Penerapan algoritma Naïve Bayes dalam mengklasifikasikan sentimen ulasan aplikasi stockbit dapat menjadi referensi dan landasan bagi penelitian selanjutnya yang berfokus pada klasifikasi teks berbasis algoritma probabilistik. Selain itu, hasil penelitian ini dapat memperkaya kajian terkait efektivitas algoritma Naïve Bayes dalam konteks fintech, yang masih jarang dibahas dalam literatur ilmiah.

2. Manfaat Praktis

Secara praktis, penelitian ini memberikan manfaat sebagai berikut:

- 1) Bagi penulis, penelitian ini merupakan sebuah eksplorasi teori-teori yang selama ini dipelajari, serta menambah wawasan, ilmu pengetahuan, dan pengalaman terhadap analisis sentimen
- 2) Bagi Universitas, sebagai tolak ukur pengetahuan mahasiswa dalam menguasai ilmu sudah dipelajari dan sebagai referensi untuk penelitian selanjutnya.
- 3) Bagi pembaca, memberikan informasi mengenai sentimen terhadap kepuasan pengguna aplikasi Stockbit dan bermanfaat untuk referensi penelitian analisis sentimen.

Bab 2 Studi Literatur

2.1 Tinjauan Penelitian Terdahulu

Penelitian ini mengacu pada penelitian terdahulu yang relevan dengan penelitian yang akan dilakukan. Berikut beberapa hasil penelitian yang relevan atau berhubungan dengan penelitian yang dijadikan bahan kajian bagi peneliti antara lain:

Menurut (Rahmawati, Rika Fitriani, & Yunizar Pratama Yusuf, 2024) yang mengimplementasikan algoritma tersebut untuk mengklasifikasikan opini pengguna terhadap sebuah aplikasi mobile perbankan berbasis Android. Penelitian tersebut memberikan performa klasifikasi yang tinggi, dengan akurasi sebesar 82%, precision sebesar 84%, *recall* sebesar 82%, dan F1-score sebesar 81%. Hasil penelitian tersebut juga menunjukkan bahwa mayoritas ulasan yang dikumpulkan dari platform distribusi aplikasi bersifat negatif. Penelitian tersebut juga membandingkan kinerja algoritma lain seperti Random Forest dan Logistic Regression, dan ditemukan bahwa Naïve Bayes memiliki keunggulan dari sisi efisiensi waktu komputasi serta akurasi klasifikasi yang lebih baik pada dataset berukuran menengah. Tahapan penting dalam penelitian tersebut mencakup pemrosesan data, seperti pembersihan teks, penghilangan kata-kata tidak penting, dan stemming, sebelum data digunakan dalam proses klasifikasi.

(Nurdin, Alexandri, Sumadinata, & Arifianti, 2023) melakukan penelitian mengenai analisis sentimen dalam konteks teknologi finansial telah berkembang pesat seiring meningkatnya penggunaan platform digital untuk layanan keuangan. Beberapa studi sebelumnya telah memanfaatkan algoritma klasifikasi seperti Naïve Bayes dan Support Vector Machine (SVM) untuk mengkaji opini publik terhadap layanan fintech yang berbeda. Salah satu penelitian membandingkan efektivitas dua pendekatan fintech, yaitu sistem tradisional dan teknologi baru, melalui analisis opini yang dikumpulkan dari berbagai sumber daring. Penelitian tersebut menggunakan algoritma SVM dan NB untuk mengklasifikasikan opini ke dalam sentimen positif atau negatif. Hasilnya menunjukkan bahwa algoritma SVM memiliki tingkat akurasi rata-rata sebesar 87,32%, sementara Naïve Bayes berada pada angka 81,56%. Selain itu, ditemukan bahwa mayoritas opini positif sekitar 71% lebih condong kepada penggunaan teknologi fintech baru. Namun demikian, tidak sedikit pula opini negatif terhadap teknologi tersebut yang disebabkan oleh isu-isu seperti keamanan data dan kurangnya pemahaman pengguna terhadap fitur yang disediakan.

Penelitian oleh (Ksatria, Yuneфри, & FC, 2023) dengan penelitian mengenai analisis sentimen dalam konteks My Pertamina menggunakan K-Nearest Neighbor dan Naïve Bayes membahas klasifikasi sentimen ulasan pengguna aplikasi My Pertamina yang diambil dari Google Playstore. Penelitian ini menggunakan dua algoritma yaitu K-Nearest Neighbor dan Naïve Bayes, dengan penerapan langsung menggunakan bahasa pemrograman Python di platform Google Collaboratory. Hasil klasifikasi menunjukkan bahwa metode K-Nearest Neighbor menghasilkan 117 ulasan positif dan 1099 ulasan negatif, sementara metode Naïve Bayes menghasilkan 84 ulasan positif dan 1132 ulasan negatif. Selain itu visualisasi hasil klasifikasi ditampilkan dalam bentuk *wordcloud* untuk sentimen positif dan negatif. Berdasarkan evaluasi performa, metode K-Nearest Neighbor menghasilkan tingkat akurasi sebesar 85,97%,

sedangkan Naïve Bayes hanya mencapai 70,73%. Dengan demikian, disimpulkan bahwa metode K-Nearest Neighbor lebih unggul dalam melakukan klasifikasi teks ulasan dibandingkan metode Naïve Bayes pada kasus ini.

Menurut (Nevita Cahaya Ramadani, 2024) dengan penelitian dengan konteks ulasan aplikasi. data ulasan diperoleh dari platform Google Playstore aplikasi Mobile Legend kecenderungan ulasan yang disampaikan mengandung komentar netral. Dari 3989 data pengujian, terdapat 1265 komentar dengan sentimen netral, kemudian 1115 komentar dengan sentimen positif, dan 1204 komentar dengan sentimen negatif. Disini dapat disimpulkan bahwa pengguna Mobile Legend di Indonesia bersikap netral terhadap aplikasi. Berdasarkan hasil akurasi algoritma SVM sebesar 87% lebih tinggi dibandingkan akurasi Algoritma *Random Forest*, Naïve Bayes, *Decision Tree* dan Algoritma *Logistic Regression* maka performa algoritma SVM sudah cukup baik.

Disisi lain, (Ramadhan et al., 2024) melakukan analisis terhadap 2000 ulasan pengguna aplikasi DANA di Google Play Store. Penelitian ini memanfaatkan algoritma Naïve Bayes dan menunjukkan akurasi sebesar 86% dengan dominasi sentimen positif, serta mengidentifikasi bahwa keluhan pengguna dalam penggunaan aplikasi meliputi kegagalan transaksi dan masalah login aplikasi.

Berdasarkan beberapa penelitian tersebut. Dapat diambil kesimpulan bahwa penelitian tersebut membahas kasus sejenis, yaitu membahas tentang penggunaan algoritma naïve bayes dalam memberikan hasil pada klasifikasi sentimen. Secara keseluruhan, Naïve Bayes mampu memberikan hasil yang cukup akurat dengan proses komputasi yang sederhana dan efisien. Namun, terdapat beberapa catatan penting seperti kebutuhan pengolahan data lebih lanjut, terutama terkait ketidakseimbangan data serta penanganan sentimen netral. Berdasarkan penelitian-penelitian tersebut, penggunaan Naïve Bayes pada analisis sentimen ulasan aplikasi Stockbit diharapkan mampu memberikan hasil klasifikasi yang baik serta membantu pengembang memahami persepsi pengguna aplikasi.

2.2 Kajian Teoritis

Berisi teori yang relevan, penjelasan secara teoritis masing-masing variabel penelitian dan landasan teori.

2.2.1 Analisis Sentimen

Analisis sentimen merupakan salah satu contoh dari bidang *Natural Language Processing* (NLP) yang paling populer. *Natural Language Processing* (NLP) adalah bidang ilmiah yang membahas tentang bagaimana caranya agar komputer bisa bekerja dan berpikir seperti manusia. *Natural Language Processing* (NLP) merupakan bagian dari *Artificial Intelligence* atau kecerdasan buatan. Dalam perkembangan data mining, *Artificial Intelligence* (AI) merupakan salah satu dari empat cabang ilmu data mining, yaitu statistika, *database*, dan pencarian informasi. Dalam penerapannya, *Artificial Intelligence* (AI) juga memerlukan *machine learning* untuk menggantikan manusia dalam mengambil keputusan. *Machine learning* tidak mempunyai perasaan seperti manusia sehingga keputusan yang diambil berdasarkan data yang sudah diolah (Irwansyah Saputra, 2022).

Analisis sentimen merupakan proses untuk mengidentifikasi dan mengkategorikan opini atau perasaan pengguna terhadap suatu topik menjadi kelas positif, negatif, atau netral. Data dari

media sosial, seperti Twitter sering dimanfaatkan untuk keperluan ini karena menyediakan banyak opini yang relevan terhadap berbagai isu atau kebijakan (Ansori, Fahmi, & Holle, 2022).

Keunggulan analisis sentimen ini adalah menghemat waktu dan tenaga dalam melakukan penelitian dengan jumlah data yang besar. Berikut contoh penerapan analisis sentimen:

1. Bidang bisnis, misalnya mengetahui bagaimana reputasi merek produk baru di masyarakat sehingga dapat meningkatkan merek produk tersebut.
2. Bidang politik, misalnya mengetahui popularitas seorang tokoh sehingga dapat memberi pengetahuan mengenai tokoh tersebut.
3. Program, misalnya vaksinasi, covid 19, pilpres, dan lain-lain sehingga dapat memperbaiki kebijakan pemerintah tersebut.

Secara umum, analisis sentimen terbagi menjadi lima langkah yaitu crawling data, pre-processing, feature selection, *classification*, dan evaluation. Analisis sentimen dapat mengubah data tidak beraturan menjadi data yang tersusun rapi. Manfaat adanya analisis sentimen yaitu sebagai evaluasi dan ide pada berbagai bidang. Analisis sentimen dapat menganalisis suatu kejadian, pernyataan, dan komentar yang kontroversi. Hasil dari analisis sentimen juga dapat menjadi sebuah gambaran bagi perusahaan, *public figure*, dan pemerintahan untuk menentukan langkah selanjutnya.

Terdapat beberapa jenis analisis sentimen yaitu *emotion detection*, *aspect based sentiment analysis*, dan *fine grand sentiment analysis*. *Fine sentiment analysis* adalah jenis analisis yang memiliki penilaian spesifik dan biasa digunakan pada bidang *e-commerce*. *Emoticon detection* adalah jenis analisis yang bertujuan untuk mengetahui emosi yang ada pada pesan misalnya emosi bahagia, sedih, marah, dan lain-lain. *Aspect-based sentiment analysis* merupakan jenis analisis untuk mengetahui aspek yang berpengaruh dan penilaian dari pelanggan (Arviana, 2021).

2.2.2 Machine Learning dan Natural Language Processing

Machine learning adalah cabang dari kecerdasan buatan (*Artificial Intelligence*) yang memungkinkan sistem untuk belajar dari data dan meningkatkan kinerjanya tanpa pemrograman eksplisit. *Machine Learning* digunakan dalam berbagai aplikasi seperti pencarian internet, deteksi spam, rekomendasi personal, perangkat lunak perbankan, dan pengenalan gambar. Prosesnya melibatkan pengumpulan data, persiapan data, pelatihan model, evaluasi, dan penerapan model untuk membuat prediksi atau keputusan berdasarkan input baru.

Natural Language Processing (NLP) adalah bidang dalam AI yang memungkinkan komputer untuk memahami, menafsirkan, dan merespon bahasa manusia. NLP digunakan dalam berbagai aplikasi seperti chatbot, analisis sentimen, dan penerjemahan mesin.

2.2.3 Google Play Store

Google Play Store adalah layanan yang disediakan oleh google yang menyediakan berbagai konten digital, diantaranya yaitu aplikasi, permainan, dan lain sebagainya. Google play store dapat diakses melalui android, situs web dan google TV. Berdasarkan Data *Mobile Operating System Market Share in Indonesia*, terdapat 89,77% pengguna Android dan 10,12% pengguna IOS yang berarti pengguna Marketplace Play Store lebih banyak dibanding dengan APP Store,, pada Google Play Store terdapat kolom penilaian dan ulasan para pengguna untuk aplikasi yang

telah tersedia (Aida Sapitri & Fikry, 2023).

Diluncurkan pertama kali pada tanggal 6 maret 2012, Google Play Store mengintegrasikan layanan Android Market, Google Music, dan Google eBookstore menjadi satu layanan terpadu yang memudahkan pengguna dalam mengakses beragam kebutuhan digital mereka. Melalui Google Play Store, pengguna dapat mengunduh aplikasi baik yang bersifat gratis maupun berbayar, memberikan ulasan (*review*), serta memberikan rating terhadap aplikasi yang diunduh.

Selain itu, fitur ulasan pengguna (*user review*) pada Google Play Store memegang peran penting dalam memberikan informasi mengenai kepuasan dan pengalaman pengguna dalam sebuah aplikasi. Ulasan ini sering digunakan oleh pengembang untuk mengevaluasi kualitas aplikasi dan melakukan perbaikan berkelanjutan. Ulasan pengguna di Google Play Store sangat berguna dalam proses analisis sentimen, karena data yang tersedia bersifat *real-time* dan mencerminkan opini langsung dari pengguna. Oleh karena itu, Google play Store merupakan salah satu sumber data yang banyak digunakan dalam penelitian analisis sentimen terhadap aplikasi.

2.2.4 Ulasan Pengguna

Ulasan di Google Play adalah fitur yang memungkinkan pengguna untuk memberikan pendapat, pengalaman, dan evaluasi tentang aplikasi yang mereka gunakan. Ulasan ini bermanfaat bagi pengembang dan pengguna lain, tetapi juga memiliki tantangan seperti jumlah, variasi bahasa, dan ketidaksesuaian rating dan teks (Daryfayi, Daulay, & Asror, 2020). Ulasan pada google play store mengandung opini dari para pengguna sebuah aplikasi sebagai pertimbangan sebelum memutuskan untuk menggunakan aplikasi.

2.2.5 Stockbit

Stockbit adalah sebuah platform investasi yang menyediakan layanan diskusi, analisis. Serta perdagangan saham secara daring. Awalnya, stockbit berfungsi sebagai komunitas daring bagi para investor untuk berbagi informasi dan analisis mengenai saham. Namun dengan perkembangan teknologi dan kebutuhan pasar, stockbit berkembang menjadi platform yang mengintegrasikan berbagai layanan investasi, termasuk riset saham, data keuangan, dan fitur perdagangan saham yang bekerja sama dengan perusahaan sekuritas.

Relevan karena banyaknya opini dan diskusi yang terjadi di dalam komunitasnya. Keunggulan dalam konteks analisis sentimen, Stockbit menjadi sumber data dalam mendukung analisis sentimen meliputi data real-time, interaksi pengguna, dan aksesibilitas data.

2.2.6 Algoritma Naïve Bayes

1) Definisi

Naïve Bayes adalah salah satu algoritma klasifikasi berbasis probabilistik yang digunakan secara luas dalam pemrosesan bahasa alami (Natural Language Processing), terutama dalam klasifikasi teks seperti analisis sentimen. Algoritma ini didasarkan pada Teorema Bayes dengan asumsi bahwa fitur-fitur (kata dalam teks) saling independen satu sama lain dalam menentukan kelas target (*class label*) meskipun dalam kenyataannya sering kali tidak demikian.

(Zhang, 2004)

2) Teorema Bayes

Dasar dari algoritma Naïve Bayes adalah Teorema Bayes, yang dinyatakan sebagai berikut:

$$P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)} \quad 2.1$$

Tabel 2. 1 Komponen teorema bayes

Rumus	Keterangan
$P(C X) =$	Probabilitas kelas C diberikan data X (posterior probability)
$P(X C) =$	Probabilitas data X diberikan kelas C
$P(C) =$	Probabilitas awal kelas C (prior probability)
$P(X) =$	Probabilitas data X secara keseluruhan

3) Asumsi Kemandirian (Naivitas)

Ciri khas Naïve Bayes adalah anggapan bahwa setiap fitur (kata) saling independen terhadap fitur lainnya dalam konteks kelas tertentu. Misalnya, jika terdapat kata “cepat” dan “mudah” dalam ulasan maka Naïve Bayes menganggap kedua tersebut tidak saling bergantung satu sama lain, yang menyederhanakan perhitungan probabilitas gabungan. Meskipun asumsi ini sederhana dan seringkali tidak benar dalam praktik, hasil klasifikasinya tetap baik karena bias-variance tradeoff (Domingos, 1997).

4) Jenis-Jenis Naïve Bayes

- Multinomial Naïve Bayes : umum digunakan untuk teks, menghitung probabilitas berdasarkan frekuensi kata
- Bernoulli Naïve Bayes : umum digunakan untuk teks, menghitung probabilitas berdasarkan frekuensi kata

2.2.7 Web Scraping

Web scraping merupakan teknik otomatisasi untuk mengekstrak data dari halaman web yang tidak menyediakan akses API atau antarmuka data langsung. Proses ini dilakukan dengan melakukan permintaan ke situs web, mengambil kontennya (umumnya dalam format HTML), dan kemudian mengekstrak informasi yang relevan berdasarkan struktur DOM-nya.

Menurut (Firza & Bakiu, 2025), web scraping telah menjadi pendekatan dominan dalam pengambilan data tidak terstruktur dari internet, terutama untuk data berbasis teks seperti ulasan pengguna, komentar, atau deskripsi produk. Penggunaan teknik ini memungkinkan peneliti dan pengembang memperoleh data dalam jumlah besar untuk keperluan analisis seperti analisis sentimen, pemodelan perilaku, dan peramalan pasar.

2.2.8 Google Colab

Google Colab adalah platform cloud computing yang mirip dengan Jupyter Notebook dan Google Research. Google Colab memungkinkan pengguna untuk menulis dan mengeksekusi kode program Python secara acak hanya dengan menggunakan browser web. Sehingga dapat memanfaatkan server untuk memproses dan menganalisis data dengan cepat

2.2.9 Python

Python adalah bahasa pemrograman populer yang bisa digunakan untuk berbagai tujuan, seperti membangun website, menganalisis data dalam *data science*, proses *scripting*, hingga pembuatan game. Python merupakan bahasa pemrograman open source, sehingga Anda bisa menggunakannya secara gratis dan bahkan ikut berkontribusi dalam pengembangannya. Banyak programmer sepakat bahwa Python adalah bahasa pemrograman yang interpretatif dan serbaguna. Sintaksnya yang mudah dibaca dan dipahami membuatnya cocok dipelajari oleh pemula (N.H & Azhar, 2024)

2.2.10 Scikit Learn

Scikit-learn adalah library bahasa pemrograman Python yang dirancang untuk memudahkan implementasi *machine learning*. Library ini dikembangkan dengan menggunakan bahasa Python, Cython, C dan C++ sehingga menghasilkan performa yang optimal dalam pemrosesan data dan penghitungan matematis. Scikit-learn menyediakan berbagai algoritma *machine learning* seperti klasifikasi, regresi, klusterisasi dan reduksi dimensi yang biasa digunakan dalam aplikasi dunia industri dan akademik. Library ini dibangun diatas modul NumPy dan SciPy, yang membuat proses komputasi nya menjadi lebih efisien. Meski demikian, scikit-learn kurang direkomendasikan untuk proses data yang berukuran sangat besar karena keterbatasan dalam scalling (Fahmi, 2023).

2.2.11 Imbalanced-learn

Imbalanced-learn adalah sebuah library Python yang dirancang khusus untuk menangani masalah data yang tidak seimbang (*imbalanced data*) dalam machine learning. Data tidak seimbang terjadi ketika rasio jumlah data antara kelas mayoritas dan kelas minoritas sangat berbeda, sehingga model cenderung bias terhadap kelas minoritas dan kurang akurat dalam memprediksi kelas minoritas. *Imbalanced-learn* menyediakan berbagai teknik untuk menyeimbangkan data seperti oversampling, undersampling, dan kombinasi keduanya, dengan metode populer seperti SMOTE (*Synthetic Over-Sampling Technique*) yang mensintesis data baru untuk kelas minoritas agar proporsinya setara dengan kelas mayoritas. Penggunaan imbalanced-learn membantu meningkatkan performa model terutama dalam kasus klasifikasi yang menghadapi distribusi kelas yang tidak seimbang (Wijaya, 2024)

2.2.12 API Google Play-Scraper

API Google-Play-Scraper adalah metode untuk mengambil data dari Google Play Store tanpa dependensi eksternal menggunakan bahasa pemrograman python. Data yang diambil dapat berupa informasi aplikasi seperti judul aplikasi, developer url, kategori aplikasi,

keseluruhan rating dan review, deskripsi, thumbnail, rating konten dan screenshot aplikasi. Selain informasi dari aplikasi API google play-scraper juga dapat mengambil data ulasan pengguna seperti nama, foto, rating, tanggal, comment likes, dan comment (Hakim, n.d.)

Berikut ini langkah-langkah dalam mengambil data ulasan di Google Play Store menggunakan API Google-Play-Scraper (Larasati, Ratnawati, & Hanggara, 2022):

- 1) Melakukan install API Google-Play-Scraper dengan menjalankan ‘pip install google play scraper’.
- 2) Melakukan import packages ‘from google_play_scraper import app’ ‘import pandas as pd’ ‘import numpy as np’.
- 3) Mengambil aplikasi ID yang terdapat di url seperti ‘id.dana’
- 4) Melakukan scraping ulasan
- 5) Memasukan hasil scraping ke dalam data frame

2.2.13 Labelling

Dalam analisis sentimen berbasis teks, khususnya pada ulasan aplikasi di platform seperti Google Play Store, proses pelabelan (*labelling*) merupakan langkah penting untuk membentuk dataset yang dapat digunakan oleh model pembelajaran mesin. Salah satu pendekatan otomatis yang umum digunakan adalah *labelling* berdasarkan skor (*score-based labelling*).

Score-based labelling menggunakan nilai rating numerik yang diberikan pengguna biasanya dalam rentang 1 sampai 5 bintang, untuk menetapkan label sentimen dari suatu ulasan. Metode ini dikembangkan sebagai solusi efisien terhadap keterbatasan pelabelan manual yang memerlukan waktu lama. Dalam proses pelabelan data, sistem menggunakan pembagian nilai rating antara 1 sampai 5, di mana skor 1 hingga 3 dikategorikan sebagai sentimen negatif, dan skor 4 hingga 5 dikategorikan sebagai sentimen positif. Pendekatan ini memungkinkan otomatisasi labeling berdasarkan data rating pengguna (Alinda Rahmi & Wulan Dari, 2024).

2.2.14 Preprocessing

Preprocessing data merupakan tahap awal yang sangat penting dalam analisis sentimen berbasis text, khususnya untuk data ulasan aplikasi dari pengguna. Tahapan ini bertujuan untuk membersihkan data mentah dari noise dan mengubahnya menjadi format yang siap diolah oleh algoritma klasifikasi sentimen. Ulasan pengguna aplikasi umumnya mengandung berbagai elemen yang tidak relevan seperti emoji, URL, tanda baca berlebihan, hingga kata-kata yang tidak baku, sehingga perlu proses penyaringan dan normalisasi terlebih dahulu.

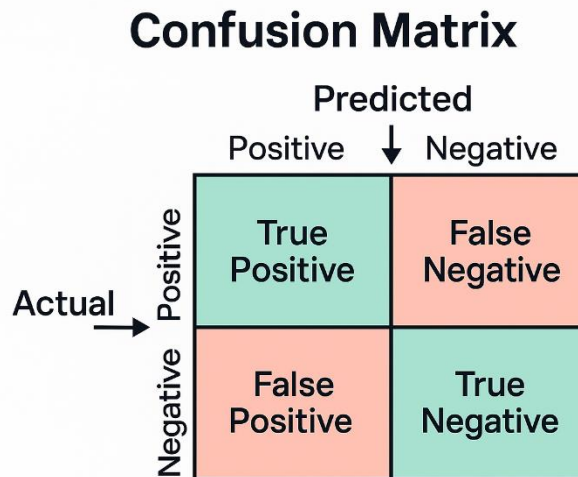
Tahapan umum dalam preprocessing mencakup beberapa langkah utama:

1. *Case Folding*, yaitu mengubah seluruh huruf menjadi huruf kecil
2. Tokenisasi, yaitu memisahkan teks ke dalam unit kata
3. *Stopword Removal* yaitu menghapus kata kata umum yang tidak memiliki nilai semantik tinggi seperti “dan”, “yang”, atau “adalah”
4. *Stemming* atau *Lemmatization*, yaitu mengembalikan kata kepada bentuk dasarnya
5. *Cleaning*, yaitu menghapus karakter spesial, angka, dan simbol yang tidak relevan.

2.2.15 Confusion Matrix

Confusion Matrix adalah tabel yang digunakan untuk memvisualisasikan kinerja model

prediksi dalam pembelajaran supervised learning. Setiap data dalam masing-masing kelas pada tabel ini menunjukkan jumlah prediksi yang dibuat untuk mengklasifikasikan kelas dengan negatif atau positif. Untuk kasus dengan dua kelas (negatif, positif), confusion matrix akan berukuran 2x2, dengan baris mewakili kelas aktual dan kolom mewakili kelas prediksi (Khalimi, n.d.).



Gambar 2. 1 Confusion matrix

a. Akurasi

Akurasi adalah matrik evaluasi yang mengukur seberapa sering model membuat prediksi yang benar dari semua prediksi yang dilakukan. Dalam klasifikasi, akurasi menunjukkan seberapa sering model memprediksi kelas yang tepat, baik positif maupun negatif. Untuk menghitung nilai akurasi, kita dapat menggunakan persamaan 1:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad 2.2$$

b. Presisi

Presisi adalah matrik evaluasi yang mengukur ketepatan model dalam membuat prediksi benar untuk kelas positif dari seluruh prediksi positif yang dibuat. Dalam klasifikasi, presisi menunjukkan seberapa sering model secara akurat memprediksi kelas positif di antara semua prediksi positif yang dihasilkan. Untuk menghitung nilai presisi, pertama-tama hitung presisi untuk setiap kelas, kemudian jumlahkan nilai-nilai tersebut dan cari rata-ratanya, Rumusnya bisa menggunakan persamaan 2:

$$Precision = \frac{Precision A + B + C}{TP + FP} \quad 2.3$$

c. Recall

Recall adalah matrik evaluasi yang menunjukkan kemampuan model dalam mengidentifikasi kelas positif secara akurat. Untuk menghitung nilai *Recall*, kita dapat menggunakan persamaan:

$$Recall = \frac{TP}{TP + FN} \quad 2.4$$

d. F1-Score

F1-Score adalah matrik evaluasi yang menggambarkan keseimbangan antara Presisi (*Precision*) dan Sensitivitas (*Recall*). Untuk menghitung nilai *F1-Score*, kita dapat menggunakan persamaan.

$$F1\ Score = \frac{Recall \times Precision}{Recall + Precision} \quad 2.5$$

2.2.16 Imbalance Dataset

Imbalanced dataset atau data tak seimbang adalah kondisi dalam sebuah dataset di mana distribusi kelas data tidak merata, yaitu jumlah sampel dalam satu kelas jauh lebih sedikit dibanding kelas lainnya. Dalam konteks klasifikasi machine learning, klasifikasi yang dihadapkan pada data tidak seimbang akan bias terhadap kelas mayoritas dan kurang mampu mengenali kelas minoritas. Kelas dengan jumlah sampel sedikit disebut kelas minoritas. Kelas dengan jumlah sampel lebih banyak disebut kelas mayoritas. Masalah ini umum terjadi pada berbagai aplikasi nyata seperti deteksi penipuan, diagnosis medis dan klasifikasi teks.

2.2.17 Teori SMOTE (Synthetic Minority Over-sampling Technique)

SMOTE (*Synthetic Minority Over-sampling Technique*) adalah salah satu teknik oversampling yang digunakan untuk mengatasi permasalahan *imbalanced dataset*, yaitu kondisi dimana distribusi kelas dalam data pelatihan tidak seimbang. Masalah ini umum terjadi dalam aplikasi klasifikasi seperti deteksi penipuan, diagnosis penyakit, dan analisis sentimen, di mana kelas minoritas—yang seringkali lebih penting—terwakili dalam jumlah data yang sangat sedikit dibandingkan kelas mayoritas (Chawla, Kovács, Tinya, Németh, & Ódor, 2002).

Berbeda dari teknik oversampling konvensional yang hanya melakukan duplikasi data kelas minoritas, SMOTE menghasilkan sampel sintesis baru dengan interpolasi antara data minoritas yang ada dan tetangga terdekatnya (*k-nearest neighbors*). Teknik ini bertujuan untuk meningkatkan kemampuan generalisasi model tanpa menimbulkan overfitting akibat pengulangan data yang sama.

2.1.1 Teori ADASYN (Adaptive Synthetic Sampling)

ADASYN (*Adaptive Synthetic Sampling*) adalah sebuah metode oversampling yang dirancang untuk mengatasi masalah ketidakseimbangan kelas (*imbalanced data*) dalam pembelajaran mesin. Metode ini dikembangkan oleh (Haibo He, Yang Bai, Garcia, & Shutao Li,

2008) untuk secara otomatis menyesuaikan distribusi data minoritas melalui pembangkitan sampel sintetis secara adaptif berdasarkan tingkat kesulitan klasifikasi masing-masing. Prinsip utama ADASYN adalah bahwa bukan semua titik minoritas diperlakukan sama, melainkan hanya titik-titik yang lebih “sulit dipelajari”—yaitu yang dikelilingi oleh banyak tetangga dari kelas mayoritas—yang menjadi fokus pembangkitan data sintetis. Hal ini membantu model pembelajaran untuk lebih fokus pada daerah batas keputusan yang kritis.

2.1.2 Streamlit

Streamlit adalah *framework* Python *open-source* yang digunakan untuk membangun aplikasi web interaktif, khususnya untuk proyek data science dan machine learning (Streamlit Inc., 2021). Kelebihan Streamlit antara lain: tidak memerlukan penguasaan HTML/CSS/JavaScript, mendukung integrasi langsung dengan library Python seperti Pandas, NumPy, Scikit-learn, dan Matplotlib, Menyediakan antarmuka interaktif secara real-time.

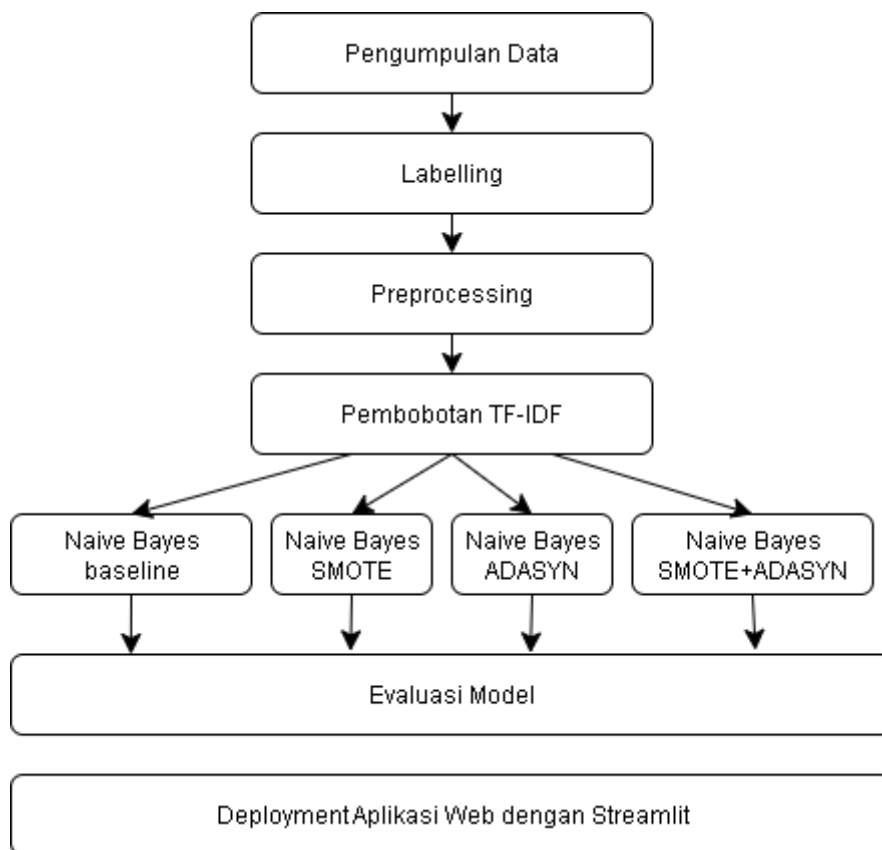
Penggunaan Streamlit untuk analisis sentimen memungkinkan model *machine learning* dipublikasikan secara cepat sehingga pengguna dapat mengunggah teks, memprosesnya, dan melihat hasil prediksi sentimen secara instan. Dalam konteks analisis sentimen, Streamlit memungkinkan model machine learning disajikan dalam bentuk aplikasi web yang memfasilitasi pengguna untuk mengunggah teks, memproses data, dan melihat hasil prediksi sentimen secara instan. Hal ini sangat berguna untuk aplikasi yang memerlukan antarmuka interaktif untuk pemantauan sentimen secara real-time dan kemudahan akses oleh pengguna yang bukan ahli teknis.

Dengan demikian, Streamlit merupakan pilihan ideal untuk membangun aplikasi berbasis machine learning dan analisis data interaktif secara cepat dan efisien, tanpa perlu keahlian pengembangan web yang mendalam.

Bab 3 Metode Penelitian

3.1 Prosedur Penelitian Eksperimental

Dalam penelitian ini, prosedur penelitian yang dilakukan melalui beberapa tahapan yaitu: pengumpulan data ulasan aplikasi Stockbit dari Google Play Store, selanjutnya tahap Pelabelan, lanjut ke tahap *Preprocessing*, pembobotan TF-IDF, Klasifikasi, dan Evaluasi Model.



Gambar 3. 1 Flowchart prosedur penelitian

3.1.1 Pengumpulan Data

Tahap pertama dalam sentimen analisis ini adalah pengumpulan data. Pada tahap ini, ulasan akan ditarik dari google play store menggunakan bahasa pemrograman Python sehingga data yang terkumpul akan menjadi sebuah dataset yang akan diberi label.(Maarif & Setiyawati, 2024).

Data ulasan yang berhasil dikumpulkan dari bulan januari 2024 hingga bulan mei 2025 sebanyak 2000 data, kemudian disimpan dalam format dokumen *comma separated values* (CSV) atau format dokumen excel (XLSX). Setelah data ulasan dikumpulkan, dilakukan penyaringan sehingga hanya menyisakan kolom ulasan dan skor, itu dilakukan karena dalam penelitian ini

hanya kolom ulasan dan skor yang dibutuhkan, selain itu dilakukan juga perurutan ulasan berdasarkan tanggal pembuatan.

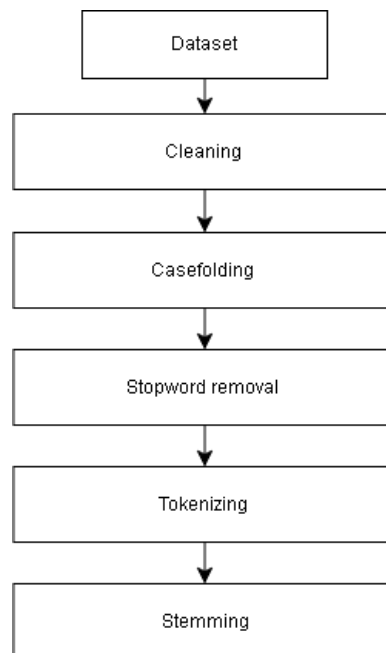


Gambar 3. 2 Flowchart pengumpulan data

Kumpulan ulasan diperoleh menggunakan teknik *scraping* data dari API yang telah disediakan oleh Google Play Store dan dimuat menjadi *python library* bernama *google-play-scraper*. Total jumlah ulasan yang penulis dapatkan adalah sebanyak 2000 ulasan (sapurta et al., 2020).

3.1.2 Tahap Preprocessing

Tahap kedua yaitu *preprocessing*, data yang sudah dikumpulkan akan diolah terlebih dahulu sebelum dilakukan klasifikasi. Pada tahap ini dataset akan dilakukan tahap preprocessing seperti pada flowchart Gambar 3.3.



Gambar 3. 3 Flowchart tahapan preprocessing

a) Cleaning

Cleaning adalah proses untuk mempersiapkan data mentah (*raw data*) agar bisa digunakan untuk melakukan analisis sentimen dengan akurat dan efektif. Tahap *cleaning* bertujuan untuk membersihkan data dari nilai yang tidak relevan, tidak valid, atau tidak diperlukan.

Pada tahap *cleaning* akan dilakukan proses untuk penghilangan tanda baca dan karakter yang tidak diperlukan seperti tanda titik, koma, tanya, seru, menghapus html, url, hashtag, mention, emoji serta menghapus karakter yang tidak relevan.

b) Case folding

Case folding adalah proses mengubah huruf kapital menjadi huruf kecil dalam teks. Dalam analisis sentimen, case folding dapat membantu mengurangi variasi kata yang sebenarnya memiliki arti yang sama, tetapi ditulis dalam huruf kapital atau kecil yang berbeda. Contohnya kata “good” dan “GOOD” memiliki arti yang sama, tetapi jika tidak dilakukan Casefolding, kedua kata ini akan dianggap berbeda dan dianalisis secara terpisah dalam proses analisis sentimen

c) Stopword Removal

Stopword removal adalah proses menghapus kata-kata yang umum dan sering muncul dalam teks, seperti kata depan, kata imbuhan, kata ganti, kata penghubung, dan kata-kata lainnya yang tidak memberikan kontribusi signifikan dalam pemahaman makna teks. Kata-kata tersebut umumnya adalah kata penghubung seperti “yang”, “dan”, “di”, “dari”, dan kata-kata lain yang tidak terlalu penting dalam analisis teks. Menghapus kata-kata yang kurang memiliki makna yang berarti seperti kata : dan, saya, atau.

Dalam analisis sentimen, penghapusan stopwords dapat membantu meningkatkan akurasi karena stopwords tidak memberikan informasi yang berguna dalam menentukan sentimen atau makna dalam suatu kalimat. Dengan menghapus stopwords, kita fokus pada kata-kata kunci yang lebih penting dalam teks untuk menentukan sentimen.

d) Tokenizing

Tokenizing adalah proses memecah teks menjadi unit-unit yang lebih kecil yang disebut “token”. Token ini dapat berupa kata, frasa, atau tanda baca. Pada langkah ini, kalimat dipisahkan menjadi potongan-potongan sebelum dianalisis lebih lanjut. Dalam analisis sentimen, tokenizing digunakan untuk mempersiapkan data teks untuk dapat diolah lebih lanjut. Data teks yang belum dipecah menjadi token-token tidak dapat diolah dalam analisis sentimen karena perlu memisahkan kata-kata dan frasa-frasa yang terdapat dalam teks menjadi unit-unit yang dapat dianalisis.

e) Stemming

Stemming adalah proses dalam *text preprocessing* untuk mengubah kata-kata menjadi kata dasar (*root word*) dengan cara memotong akhiran kata. Stemming mengubah kata-kata

yang memiliki imbuhan menjadi sebuah kata dasar aslinya. Tujuan stemming untuk mengurangi variasi kata yang muncul serta mempertahankan makna kata

f) *Word Cloud*

Word Cloud adalah representasi visual dari kumpulan kata yang menunjukkan frekuensi atau tingkat kepentingan kata-kata tertentu dalam sebuah dokumen atau korpus teks. Kata-kata yang muncul lebih sering ditampilkan dengan ukuran huruf yang lebih besar, sedangkan kata-kata yang jarang muncul ditampilkan dengan ukuran lebih kecil.

3.1.3 Pembobotan TF-IDF

Pembobotan TF-IDF (Term Frequency Inverse Document Frequency) adalah metode untuk pembobotan teks yang digunakan untuk mengevaluasi seberapa penting suatu kata dalam sebuah dokumen relatif terhadap kumpulan dokumen. Teknik ini merupakan salah satu pendekatan paling populer dalam bidang *text mining* dan *Natural Language Processing* (NLP) untuk merepresentasikan teks dalam bentuk numerik agar dapat diproses oleh algoritma pembelajaran mesin. TF-IDF terdiri dari dua komponen utama:

1. Term Frequency (TF): Mengukur seberapa sering suatu kata muncul dalam dokumen tertentu. Semakin sering suatu kata muncul, semakin tinggi nilai TF-nya. Namun, kata yang terlalu sering muncul tidak selalu informatif karena bisa jadi hanya kata umum (seperti “dan”, “atau”, dll).
2. Inverse Document Frequency (IDF): Mengukur sejauh mana kata yang memiliki sifat unik di seluruh dokumen. Jika sebuah kata muncul di banyak dokumen, maka IDF-nya rendah karena kata tersebut kurang informatif. Sebaliknya, jika hanya muncul di sedikit dokumen, maka IDF-nya tinggi, menandakan bahwa kata tersebut memiliki nilai yang lebih tinggi dalam membedakan dokumen.
3. Rumus umum TF-IDF adalah :

$$TF - IDF(t, d) = TF(t, d) \times \log \log \left(\frac{N}{DF(t)} \right) \quad 3.1$$

Tabel 3. 1 Simbol perhitungan TF-IDF

Simbol	Keterangan
t	term atau kata
d	dokumen tertentu
N	jumlah total dokumen
DF(t)	jumlah dokumen yang mengandung t

3.1.4 Teknik Resampling

Pada penelitian ini, digunakan dua metode resampling untuk menangani ketidakseimbangan data (*imbalanced data*) yang ditemukan pada distribusi kelas sentimen. Ketidakseimbangan ini dapat menyebabkan model klasifikasi bias terhadap kelas mayoritas, sehingga diperlukan pendekatan yang mampu menyeimbangkan jumlah data antar kelas. Dua teknik yang digunakan adalah SMOTE (*Synthetic Minority Over-Sampling Technique*) dan ADASYN (*Adaptive Synthetic Sampling*)

a) SMOTE

Dalam penelitian ini, SMOTE diterapkan setelah proses ekstraksi fitur (TF-IDF) dan sebelum pelatihan model klasifikasi, guna memastikan distribusi kelas menjadi seimbang.

b) ADASYN

Dalam penelitian ini, ADASYN digunakan sebagai pembanding terhadap SMOTE untuk mengukur efektivitas dalam meningkatkan performa model Naive Bayes dan SVM pada data yang tidak seimbang.

3.1.5 Pelatihan model klasifikasi

Proses klasifikasi menggunakan Algoritma Naïve Bayes terdiri dari dua tahap: proses pelatihan (*training*) dan proses pengujian (*testing*). Pertama, dilakukan proses pelatihan untuk membangun model, kemudian dilanjutkan dengan proses pengujian yang mengacu pada probabilitas dari dataset pelatihan. Metode Naïve Bayes Classifier digunakan untuk mengelompokkan opini dengan baik, mampu mengklasifikasikan komentar menjadi positif, netral, atau negatif terhadap suatu produk atau isu yang sedang ramai dibicarakan.

3.1.6 Evaluasi Model

Model ini dievaluasi menggunakan metrik akurasi, presisi, *recall*, dan F1-score. Hasil akurasi nantinya akan menunjukkan kinerja pelatihan model sebelumnya dengan menunjukkan tingkat akurasi. Evaluasi model sangat diperlukan untuk mengetahui apakah model dapat menggeneralisasi pola dan fitur yang telah dipelajari ke data baru yang belum pernah dilihat sebelumnya, sehingga hasil akurasi dari evaluasi model dapat menentukan apakah model perlu diperbaiki ataupun tidak.

3.1.7 Deployment aplikasi web dengan Streamlit

Tahap terakhir penelitian adalah deployment, yaitu mengimplementasikan model analisis sentimen ke dalam aplikasi web menggunakan Streamlit. Streamlit dipilih karena sederhana, open source, dan mudah terintegrasi dengan Python.

Prosedurnya dilakukan dengan langkah berikut:

1. Menyimpan model hasil pelatihan Naïve Bayes dalam format .pkl.
2. Membangun antarmuka menggunakan Streamlit, mulai dari input teks, preprocessing, hingga menampilkan hasil klasifikasi sentimen.
3. Uji coba lokal dengan menjalankan streamlit run app.py untuk memastikan aplikasi berjalan sesuai harapan.

4. Deploy online ke layanan hosting (misalnya *Streamlit Cloud*) agar bisa diakses melalui browser.
5. Uji akses pengguna untuk memastikan aplikasi mudah digunakan dan hasil prediksi sesuai model.

Dengan tahap ini, penelitian tidak hanya berhenti pada model, tetapi juga menghasilkan aplikasi praktis yang dapat digunakan secara langsung oleh pengguna.

3.2 Analisa Sistem

3.2.1 Hipotesa

- 1) H_0 (Hipotesis Nol): Penerapan algoritma Naïve Bayes pada data ulasan pengguna aplikasi Stockbit tidak memberikan hasil klasifikasi sentimen yang signifikan terhadap perbedaan performa antara data tanpa resampling dan dengan resampling (SMOTE/ADASYN).
- 2) H_1 (Hipotesis Alternatif): Penerapan algoritma Naïve Bayes pada data ulasan pengguna aplikasi Stockbit memberikan hasil klasifikasi sentimen yang lebih baik secara signifikan ketika dilakukan resampling menggunakan SMOTE atau ADASYN.

3.2.2 Statistik Deskripsi

Statistik deskriptif yang dihasilkan dari data ulasan antara lain:

- 1) Jumlah total data ulasan dikumpulkan
- 2) Persentase data sentimen positif dan negatif
- 3) Penerapan metode resampling (SMOTE dan ADASYN) pada data latih.
- 4) Akurasi model berdasarkan masing masing metode resampling .

3.2.3 Analisa Sesuai Tema

Analisis dilakukan untuk memahami:

1. Tantangan dalam analisis sentimen bahasa indonesia di bidang investasi
2. Evaluasi performa awal dari empat model algoritma Naïve Bayes, yaitu:
 - a) Naïve Bayes tanpa resampling,
 - b) Naïve Bayes dengan SMOTE,
 - c) Naïve Bayes dengan ADASYN,
 - d) Naïve Bayes dengan SMOTE + ADASYN
3. Perbandingan keempat model tersebut untuk menentukan model terbaik. Model terbaik yang terpilih selanjutnya akan digunakan untuk tahap deployment aplikasi

3.3 Simulasi

Simulasi dilakukan untuk menguji efektivitas metode klasifikasi sentimen terhadap ulasan aplikasi Stockbit dengan menggunakan algoritma Naïve Bayes. Dataset yang digunakan merupakan data sekunder yang diperoleh melalui proses Scraping. Seluruh tahapan mulai dari Preprocessing, ekstraksi fitur, hingga evaluasi model dilakukan dengan menggunakan pustaka Python seperti Pandas, NLTK, dan Scikit-learn. Berikut merupakan rancangan dalam simulasi:

1. Deskripsi Dataset
 - a) Sumber data : dataset diperoleh dari ulasan aplikasi Stockbit dari Google Play Store menggunakan 'Google Play Scraper'

- b) Data ulasan diambil bulan januari 2024 hingga mei 2025 sebanyak 2000 data ulasan.
2. Tahapan preprocessing
- Contoh ulasan mentah :
- “Yang Sudah Punya Akun Bibit Tinggal Koneksikan Saja ke Akun Stockbit-nya. Fitur nya Oke, Mantap lah.”

Tabel 3. 2 Contoh tahapan preprocessing

Tahap	Hasil implementasi
Cleaning	yang sudah punya akun bibit tinggal koneksikan saja ke akun stockbit nya fitur nya oke mantap lah
Case Folding	yang sudah punya akun bibit tinggal koneksikan saja ke akun stockbit nya. fitur nya oke, mantap lah
Tokenisasi	'yang', 'sudah', 'punya', 'akun', 'bibit', 'tinggal', 'koneksikan', 'saja', 'ke', 'akun', 'stockbit', 'nya', 'fitur', 'nya', 'oke', 'mantap', 'lah'
Stopword Removal	'punya', 'akun', 'bibit', 'tinggal', 'koneksikan', 'akun', 'stockbit', 'fitur', 'oke', 'mantap'
Normalisasi	'punya', 'akun', 'bibit', 'tinggal', 'koneksi', 'akun', 'stockbit', 'fitur', 'baik', 'bagus']
Stemming	'punya', 'akun', 'bibit', 'tinggal', 'koneksi', 'akun', 'stockbit', 'fitur', 'baik', 'bagus'

3. Transformasi data
- Transformasi data dilakukan menggunakan TF-IDF dengan bantuan TfidfVectorizer dari pustaka Scikit-learn. Proses ini mengubah data teks ulasan pengguna menjadi vektor angka yang dapat diproses oleh algoritma Naïve Bayes.

$$TF - IDF(t, d) = TF(t, d) \times \log \log \left(\frac{N}{DF(t)} \right) \quad 3.2$$

Dengan:

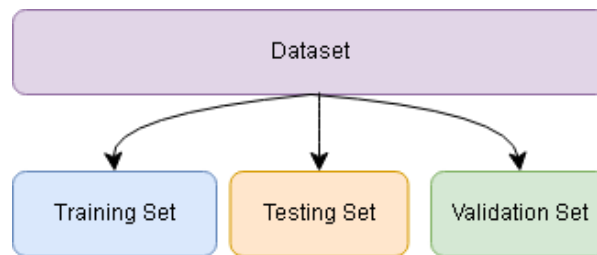
- a) t: kata atau term

- b) d: dokumen tertentu
- c) N: jumlah total dokumen
- d) $DF(t)$: jumlah dokumen yang mengandung kata t

Tabel 3. 3 Contoh perhitungan TF-IDF

Term (t)	df	Idf
Belajar	1	$\log(4/1)=0.602$
Hitung	2	$\log(4/2)=0.301$
Idf	3	$\log(4/3)=0.125$
Saya	4	$\log(4/4)=0$

4. Split Data



Gambar 3. 4 Pembagian split dataset

Split data yang dilakukan dengan membagi data menjadi 3 bagian yaitu : Data training sebesar 60-80%, Data Validation 10-20%. Dengan membagi data menjadi tiga bagian diharapkan kinerja model meningkat dalam segi kemampuan serta lebih akurat dari data sebelumnya. Rumus split data yang umum digunakan yaitu :

$$\begin{aligned}
 \text{Proporsi split} &= \frac{\text{Jumlah data Training}}{\text{Jumlah total data}} \\
 \text{Training set : } X_{\text{train}} &= X[: \text{train}_{\text{size}}] \\
 \text{Testing set : } X_{\text{train}} &= X[: \text{train}_{\text{size}}]
 \end{aligned}
 \tag{3.3}$$

4. Teknik Resampling SMOTE

Synthetic Minority Oversampling Technique (SMOTE) adalah metode *resampling* yang digunakan untuk mengatasi ketidakseimbangan kelas (*class imbalance*) pada dataset. Teknik ini bekerja dengan menghasilkan sampel sintetis (*synthetic samples*) dari kelas minoritas, bukan hanya melakukan duplikasi data yang ada.

Langkah- langkah SMOTE

a. Pencarian tetangga terdekat (Nearest Neighbors)

Untuk setiap sampel minoritas x_i , tentukan k tetangga terdekatnya menggunakan algoritma seperti k-Nearest Neighbors (k-NN).

- b. Pemilihan tetangga secara acak
Pilih salah satu tetangga terdekat x_{zi} secara acak.
- c. Interpolasi linear untuk data sintetis
Buat sampel sintetis x_{new} dengan persamaan :

$$x_{new} = x_i + \delta \times (x_{zi} - x_i) \quad 3.4$$

Tabel 3. 4 Simbol resampling SMOTE

Simbol	Keterangan
x_i	Data asli kelas minoritas
x_{zi}	Salah satu tetangga terdekat dari x_i
δ	Bilangan acak dalam rentang $[0, 1]$

- d. pengulangan
Ulangi proses ini hingga jumlah data sintetis yang diinginkan tercapai

5. ADASYN

ADASYN (Adaptive Synthetic Sampling) adalah metode oversampling adaptif yang menambah sampel sintetis terutama pada area data minoritas yang sulit dipelajari (berdekatan dengan mayoritas). Intinya mirip SMOTE (interpolasi tetangga terdekat), tetapi jumlah sampel sintetis per titik minoritas diproporsikan oleh tingkat kesulitannya. Proses ADASYN dilakukan melalui langkah-langkah berikut:

- a) Hitung rasio ketidakseimbangan data:

$$d = \frac{m_s}{m_l} \quad 3.5$$

Dengan m_s adalah jumlah sampel kelas minoritas, dan m_l adalah jumlah sampel kelas mayoritas.

- b) Tentukan jumlah sampel sintetis yang akan dibuat

$$G = (m_l - m_s) \times \beta \quad 3.6$$

Dengan $\beta \in [0,1]$, adalah parameter yang mengatur berapa banyak sampel sintetis yang dihasilkan.

- c) Hitung bobot kesulitan klasifikasi tiap sampel minoritas

Untuk setiap data minoritas x_i , hitung proporsi tetangga dari kelas mayoritas:

$$r_i = \frac{\Delta_i}{K} \quad 3.7$$

Dimana Δ_i adalah jumlah tetangga dari kelas mayoritas, dan K adalah jumlah tetangga terdekat yang digunakan.

- d) Generate data sintesis

Sampel sintesis dihasilkan secara acak mengikuti prinsip interpolasi seperti SMOTE:

$$x_{new} = x_i + \delta \times (x_{zi} - x_i) \quad 3.8$$

Dengan x_{zi} adalah tetangga terdekat yang dipilih, dan δ adalah bilangan acak uniform pada interval $[0,1]$.

6. Algoritma Naïve Bayes

Algoritma yang digunakan dalam simulasi ini adalah *Multinomial Naïve Bayes*, yang merupakan salah satu varian dari algoritma Naïve Bayes dan paling umum digunakan untuk kasus klasifikasi teks seperti analisis sentimen. Multinomial Naïve Bayes bekerja dengan menghitung probabilitas suatu dokumen termasuk ke dalam suatu kelas berdasarkan frekuensi kemunculan kata-kata dalam dokumen tersebut. Asumsinya adalah bahwa fitur (kata) dalam dokumen mengikuti distribusi *multinomial* dan bersifat saling bebas (independen). Berikut merupakan rumus dasar Naïve Bayes berdasarkan Teorema Bayes.

$$P(X) = \frac{P(C) \cdot P(C)}{P(X)} \quad 3.9$$

7. Evaluasi model

Evaluasi dilakukan untuk menggunakan fungsi *classification report* dan *confusion matrix* dari library Scikit-learn untuk mendapatkan metrik-metrik tersebut dan visualisasi hasil klasifikasi. Matrik evaluasi yang digunakan meliputi:

- 1) Accuracy: Tingkat keakuratan klasifikasi secara keseluruhan.

Rumus:

$$Akurasi = (TP + TN) / (TP + TN + FP + FN)$$

- 2) Precision: Kemampuan model dalam mengklasifikasikan data positif secara tepat.

Rumus:

Bab 5 Penutup

5.1 Kesimpulan

Berdasarkan hasil penelitian dan analisis yang telah dilakukan pada ulasan pengguna aplikasi Stockbit di Google Play Store, maka dapat disimpulkan beberapa hal sebagai berikut:

1. Algoritma Multinomial Naïve Bayes berhasil diimplementasikan untuk mengklasifikasikan sentimen ulasan menjadi dua kelas, positif dan negatif. Model baseline menunjukkan akurasi awal yang baik namun memiliki keterbatasan dalam mengenali kelas negatif pada kondisi data tidak seimbang.
2. Ketidakseimbangan data, dengan proporsi ulasan positif $>80\%$ dan negatif $\pm 20\%$, menyebabkan bias klasifikasi ke kelas mayoritas, tercermin dari rendahnya recall kelas negatif pada model awal.
3. Penerapan teknik resampling SMOTE dan ADASYN berhasil menyeimbangkan distribusi kelas serta meningkatkan performa, khususnya recall dan F1-score pada kelas negatif. ADASYN memberikan hasil adaptif dengan fokus pada sampel sulit.
4. Dari empat skenario model, NB + ADASYN dipilih sebagai model terbaik karena menunjukkan keseimbangan *precision-recall* pada kedua kelas dengan akurasi 85%, dan diimplementasikan ke aplikasi berbasis Streamlit untuk memudahkan penggunaan pada data baru.

5.2 Saran

Berdasarkan hasil penelitian dan evaluasi sistem yang telah dilakukan, beberapa saran untuk pengembangan lebih lanjut adalah sebagai berikut::

1. Penggunaan algoritma lain sebagai pembandingan sangat disarankan dalam penelitian lanjutan, seperti *Logistic Regression*, SVM, untuk mengetahui apakah ada algoritma lain yang dapat memberikan performa lebih baik terutama pada data yang *imbalanced*.
2. Peningkatan kualitas preprocessing teks dapat dilakukan lebih lanjut, misalnya dengan mempertimbangkan penggunaan lemmatization, penghapusan kata-kata ambigu, atau analisis berbasis aspek (*aspect-based sentiment analysis*) agar hasil klasifikasi lebih tajam dan kontekstual
3. Pengembangan lebih lanjut aplikasi berbasis Streamlit, misalnya dengan menambahkan fitur unggah dataset baru, filter hasil prediksi, dan penyimpanan riwayat analisis untuk mendukung pengguna non-teknis.

Referensi

- Aida Sapitri, I., & Fikry, M. (2023). Pengklasifikasian Sentimen Ulasan Aplikasi Whatsapp Pada Google Play Store Menggunakan Support Vector Machine. *Jurnal TEKINKOM*, 6(1), 1–7. <https://doi.org/10.37600/tekinkom.v6i1.773>
- Alinda Rahmi, N., & Wulan Dari, R. (2024). Implementation of Natural Language Processing (Nlp) in Consumer Sentiment Analysis of Product Comments on the Marketplace. *Jurnal Teknik Informatika (Jutif)*, 5(3), 693–701. <https://doi.org/10.52436/1.jutif.2024.5.3.1666>
- Ansori, Y., Fahmi, K., & Holle, H. (2022). Perbandingan Metode Machine Learning dalam Analisis Sentimen Twitter Comparison of Machine Learning Methods in Twitter Sentiment Analysis, 10(4), 1–6. <https://doi.org/10.26418/justin.v10i4.51784>
- Arviana, G. N. (2021). Sentiment Analysis, Teknik untuk Pahami Maksud di Balik Opini Pelanggan. Diambil 26 April 2025, dari <https://glints.com/id/lowongan/sentiment-analysis/>
- Chawla, Kovács, B., Tinya, F., Németh, C., & Ódor, P. (2002). SMOTE: Synthetic Minority Over-sampling Technique. *Ecological Applications*, 30(2), 321–357. <https://doi.org/10.1002/eap.2043>
- Daryfayi, E., Daulay, P., & Asror, I. (2020). Sentimen Analisis pada Ulasan Google Play Store Menggunakan Metode Naïve Bayes. *e-Proceeding of Engineering*, 7(2), 8400–8410.
- Fahmi, M. N. (2023). Implementasi Mechine Learning menggunakan Python Library : Scikit-Learn (Supervised dan Unsupervised Learning). *Sains Data Jurnal Studi Matematika dan Teknologi*, 1(2), 87–96. <https://doi.org/10.52620/sainsdata.v1i2.31>
- Firza, N., & Bakiu, A. (2025). Machine Learning for Quality Diagnostics : Insights into Consumer Electronics Evaluation. *MDPI Open Access Journal*.
- google play store. (2025). Stockbit - Investasi Saham & Komunitas. Diambil 24 April 2025, dari <https://play.google.com/store/apps/details?id=com.stockbit.android>
- Haibo He, Yang Bai, Garcia, E. A., & Shutao Li. (2008). ADASYN: Adaptive synthetic sampling approach for imbalanced learning. In *2008 IEEE International Joint Conference on Neural Networks (IEEE World Congress on Computational Intelligence)* (hal. 1322–1328). IEEE. <https://doi.org/10.1109/IJCNN.2008.4633969>
- Hakim, A. R. (n.d.). Scraping Website Ulasan Shopee App di Google Play. Diambil 30 April 2025, dari <https://bisa.ai/portofolio/detail/MTI3OA>
- Irwansyah Saputra, D. A. . (2022). *Machine Learning Untuk Pemula* (Maret 2022). Bandung: Informatika.
- Jazuli, A. (2024). Optimizing Aspect-Based Sentiment Analysis Using BERT for Comprehensive Analysis of Indonesian Student Feedback. *MDPI*, 1–28.
- Khalimi, A. M. (n.d.). Perhitungan Confusion Matrix Multi-Class Clasification 3x3. Diambil 30 April 2025, dari <https://www.pengalaman-edukasi.com/2020/11/menghitung-confusion-matrix-3-kelas.html>
- Ksatria, D. T., Yunefri, Y., & FC, L. L. Van. (2023). Analisis Sentimen Ulasan Pengguna Aplikasi Mypertamina pada Google Playstore Menggunakan K-Nearest Neighbor dan

- Naïve Bayes. *Teknologi Informasi &*, 2(1), 213–227. Diambil dari <https://journal.unilak.ac.id/index.php/Semaster/article/view/18526>
- Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. (2022). Analisis Sentimen Ulasan Aplikasi Dana dengan Metode Random Forest, 6(9), 4305–4313.
- Maarif, M. M., & Setiyawati, N. (2024). Analisis Sentimen Review Aplikasi LinkedIn di Google Play Store Menggunakan Support Vector Machine. *Progresif: Jurnal Ilmiah Komputer*, 20(1), 454. <https://doi.org/10.35889/progresif.v20i1.1614>
- N.H, & Azhar, M. (2024). Pengertian Python: Bahasa Pemrograman Serbaguna dan Populer. Diambil 30 April 2025, dari <https://bse.telkomuniversity.ac.id/pengertian-python-bahasa-pemrograman-serbaguna-dan-populer/>
- Nevita Cahaya Ramadani. (2024). Analisis Sentimen Untuk Mengukur Ulasan Pengguna Aplikasi Mobile Legend Menggunakan Algoritma Naive Bayes, SVM, Random Fores, Decision Tree, dan Logistic Regression Nevita. *JSI : Jurnal Sistem Informasi (E-Journal)*, 16(1), 123–138.
- Nurdin, T. A., Alexandri, M. B., Sumadinata, W., & Arifianti, R. (2023). SENTIMENT ANALYSIS OF USER PREFERENCE FOR OLD VS NEW FINTECH TECHNOLOGY USING SVM AND NB ALGORITHMS, 0. <https://doi.org/10.2478/mspe-2023-0041>
- Otoritas Jasa Keuangan. (2022). Siaran Pers: Jumlah Investor Ritel Pasar Modal Terus Meningkat. Diambil 24 April 2025, dari <https://www.ojk.go.id/id/berita-dan-kegiatan/siaran-pers/Pages/Jumlah-Investor-Ritel-Pasar-Modal-Terus-Meningkat.aspx>
- Pamuji, F. Y., & Putri, S. D. A. (2023). Komparasi Metode Smote Dan Adasyn Untuk Penanganan Data Tidak Seimbang Multiclass. *Jurnal Informatika Polinema*, 9(3), 331–338. <https://doi.org/10.33795/jip.v9i3.1330>
- Rahel Lina Simanjuntak, Theresia Romauli Siagian, Vina Anggriani, & Arnita Arnita. (2023). Analisis Sentimen Ulasan Pada Aplikasi E-Commerce Shopee Dengan Menggunakan Algoritma Naïve Bayes. *Jurnal Teknik Mesin, Elektro dan Ilmu Komputer*, 3(3), 23–39. <https://doi.org/10.55606/teknik.v3i3.2411>
- Rahmawati, I., Rika Fitriani, T., & Yunizar Pratama Yusuf, A. (2024). Analisis Sentimen Menggunakan Algoritma Naïve Bayes Pada Aplikasi m-BCA berdasarkan Ulasan Pengguna di Google Play Store. *Jurnal Riset Informatika dan Teknologi Informasi*, 1(2), 38–42. <https://doi.org/10.58776/jriti.v1i2.116>
- Ramadhan, G. R., Sugianto, C. A., Studi, P., Informatika, T., Cimahi, K., Store, G. P., & Machine, S. V. (2024). ANALISIS SENTIMEN ULASAN APLIKASI DANA DI GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA NAÏVE BAYES, 8(5), 9849–9857.
- Serlina, A., & Rahim, A. (2025). Comparative Analysis of Naïve Bayes Algorithm Performance in English and Indonesian Text Sentiment *Classification* on Duolingo Application in Playstore. *Journal of Information and Communication Technology*, 14(March), 165–171. <https://doi.org/10.34148/teknika.v14i1.1207>
- Wijaya, C. Y. (2024). 5 Teknik SMOTE untuk *Oversampling* Data yang *Imbalance*. Diambil dari <https://www.berdata.com/post/5-teknik-smote-untuk-oversampling-data-yang-imbalance>