



Research article

Ulasan Pengguna Terhadap Aplikasi SEVIMA EdLink Menggunakan Metode Machine Learning

User Reviews of the SEVIMA EdLink Application Using the Machine Learning Method

*Rizki¹, Ridho², Nur Ilham Ade Putra Ramadhan³, Muhammad Asril Fadli⁴

¹²³⁴Sistem Informasi, Teknik dan Ilmu Komputer, Universitas Islam Indragiri

Jl. Provinsi Parit 1 Tembilahan Hulu, Indragiri Hilir

email: ¹rizkirizki3941@gmail.com, ²ridhomdta@gmail.com, ³nurilhamade3@gmail.com, ⁴asrilfadli071@gmail.com

*Correspondence

ARTICLE INFO

Article history:

Received November 25, 2025

Revised November 30, 2025

Accepted May 30, 2026

Keywords:

KDD 2

EdLink 3

TF-IDF

Naïve Bayes

SVM

Logistic Regression

ABSTRACT

Penelitian ini bertujuan untuk melakukan analisis sentimen pada ulasan pengguna aplikasi EdLink menggunakan pendekatan Knowledge Discovery in Database (KDD), yang terdiri dari tahap data selection, preprocessing, transformation, data mining, dan evaluation. Sebanyak 2.509 ulasan berhasil dikumpulkan menggunakan google-play-scraper, kemudian dilakukan pembersihan teks dan pelabelan otomatis berdasarkan skor rating. Data teks ditransformasikan menggunakan metode TF-IDF sebelum dilakukan klasifikasi menggunakan tiga algoritma machine learning, yaitu Naive Bayes, Support Vector Machine (SVM), dan Logistic Regression. Hasil pengujian menggunakan tiga skema pembagian data (50–50, 70–30, dan 80–20) menunjukkan bahwa model Naive Bayes memiliki performa terbaik dengan akurasi tertinggi mencapai 0,77 pada split 80–20, disusul Logistic Regression sebesar 0,76, sedangkan SVM memiliki akurasi terendah. Temuan ini memberikan gambaran bahwa meskipun aplikasi EdLink dinilai bermanfaat, pengguna masih sering mengalami kendala teknis. Penelitian ini diharapkan dapat menjadi bahan evaluasi pengembang untuk meningkatkan kualitas aplikasi.

This research aims to conduct sentiment analysis on user reviews of the EdLink application using the Knowledge Discovery in Database (KDD) approach, which consists of data selection, preprocessing, transformation, data mining, and evaluation stages. A total of 2,509 reviews were successfully collected using the google-play-scraper, followed by text cleaning and automatic labeling based on rating scores. The text data was transformed using the TF-IDF method before classification was performed using three machine learning algorithms: Naive Bayes, Support Vector Machine (SVM), and Logistic Regression. The test results using three data splitting schemes (50-50, 70-30, and 80-20) showed that the Naive Bayes model performed best with the highest accuracy reaching 0.77 in the 80-20 split, followed by Logistic Regression at 0.76, while SVM had the lowest accuracy. These findings illustrate that although the EdLink application is considered beneficial, users still frequently experience technical obstacles. This research is expected to serve as evaluation material for developers to improve the quality of the application.

1. Pendahuluan

Perkembangan teknologi digital dalam sektor pendidikan telah mendorong lahirnya berbagai aplikasi pembelajaran yang digunakan untuk mendukung proses komunikasi, manajemen kelas, dan distribusi materi secara daring. Salah satu aplikasi yang banyak digunakan oleh mahasiswa dan dosen di Indonesia adalah SEVIMA EdLink, sebuah platform pembelajaran digital yang menyediakan fitur ruang kelas, pengumpulan tugas, absensi, dan komunikasi akademik secara terpadu. Dengan tingginya intensitas penggunaan aplikasi ini, opini pengguna yang tercermin melalui ulasan pada Google Play Store menjadi sumber informasi penting untuk mengevaluasi kualitas layanan dan pengalaman pengguna secara langsung.

Ulasan pengguna bersifat tidak terstruktur, beragam, dan jumlahnya sangat besar sehingga sulit dianalisis secara manual. Oleh karena itu, analisis sentimen menjadi pendekatan yang relevan untuk mengidentifikasi kecenderungan opini pengguna melalui pemrosesan bahasa alami. Analisis sentimen telah banyak digunakan untuk mengevaluasi kualitas aplikasi digital dan memahami kepuasan pengguna. Misalnya, penelitian oleh [1]

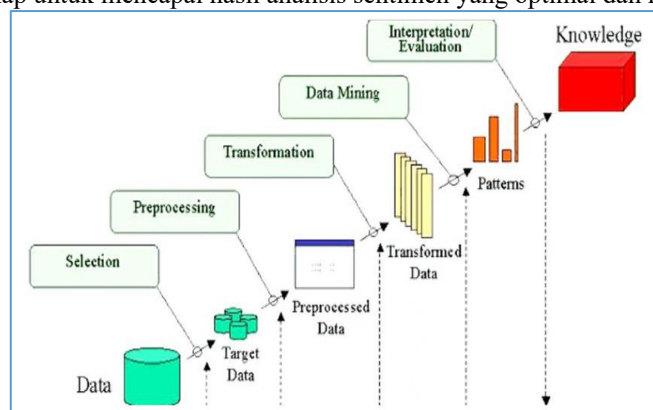
menunjukkan bahwa analisis sentimen pada ulasan aplikasi Shopee dapat memberikan gambaran yang jelas mengenai keluhan dan apresiasi pengguna, serta membantu pengembang meningkatkan layanan aplikasi e-commerce. Temuan serupa juga didukung oleh penelitian [2], yang menunjukkan bahwa metode klasifikasi seperti Naive Bayes dan Support Vector Machine efektif dalam mengelompokkan sentimen pengguna terhadap aplikasi Pluang berdasarkan ulasan di Google Play Store.

Pada ranah aplikasi pendidikan, analisis sentimen juga terbukti bermanfaat. Penelitian yang dilakukan oleh (Sukmawati et al., 2021)[3] terhadap aplikasi Glints menunjukkan bahwa ulasan pengguna dapat diolah menjadi informasi berharga untuk mengevaluasi fungsionalitas, kenyamanan penggunaan, serta permasalahan teknis yang sering dialami pengguna. Meski demikian, penelitian terkait analisis sentimen pada aplikasi pembelajaran lokal seperti EdLink masih sangat terbatas. Sebagian besar penelitian sebelumnya berfokus pada aplikasi komersial dan belum memanfaatkan dataset ulasan pendidikan dalam skala besar. Hal ini menunjukkan adanya kebutuhan untuk melakukan penelitian yang lebih dalam terhadap persepsi pengguna aplikasi EdLink sebagai representasi penggunaan aplikasi pembelajaran di lingkungan perguruan tinggi.

Berdasarkan celah tersebut, penelitian ini bertujuan untuk menganalisis sentimen pengguna aplikasi EdLink menggunakan pendekatan Knowledge Discovery in Database (KDD) yang meliputi tahapan data selection, preprocessing, transformation, data mining, dan evaluation. Penelitian ini membandingkan performa tiga algoritma klasifikasi, yaitu Naive Bayes, Support Vector Machine, dan Logistic Regression, untuk mengidentifikasi sentimen positif, negatif, dan netral dari ribuan ulasan pengguna. Hasil penelitian ini diharapkan dapat memberikan insight yang komprehensif mengenai persepsi pengguna dan dapat menjadi masukan bagi pengembang EdLink dalam meningkatkan kualitas aplikasi di masa mendatang.

2. Metode Penelitian

Metode penelitian ini menggunakan pendekatan Knowledge Discovery in Database (KDD) sebagai kerangka kerja utama untuk menggali pengetahuan dari data ulasan pengguna aplikasi EdLink. KDD dipilih karena menyediakan alur sistematis yang mencakup pengumpulan data, pembersihan, transformasi, hingga penerapan model dan evaluasi hasil[4]. Selain itu, dalam konteks analisis ulasan aplikasi mobile, studi-studi lokal menunjukkan bahwa proses yang terstruktur seperti KDD penting untuk memastikan model klasifikasi sentimen berjalan dengan baik dan dapat dipertanggungjawabkan[5]. Dengan demikian, penelitian ini menjalankan tiap tahap KDD secara bertahap untuk mencapai hasil analisis sentimen yang optimal dan relevan secara aplikasi.



Gambar 1. Proses KDD

Dengan menggunakan pendekatan ini, penelitian dapat dilakukan secara lebih terarah, mulai dari tahap awal pengolahan data hingga menghasilkan model prediktif yang berkualitas. Oleh karena itu, pada bagian berikutnya akan dijelaskan tahapan-tahapan KDD yang digunakan dalam penelitian ini.

2.1 Data Selection (Pemilihan Data)

Tahap Data Selection merupakan langkah awal yang sangat penting dalam proses KDD karena menentukan kualitas data yang akan diproses pada tahapan selanjutnya. Pada tahap ini, peneliti memilih data yang relevan dan sesuai dengan tujuan penelitian. Data yang dipilih dalam penelitian ini adalah ulasan pengguna aplikasi EdLink yang tersedia pada platform Google Play Store. Pemilihan ulasan dilakukan dengan mempertimbangkan atribut yang mendukung analisis sentimen, seperti teks ulasan (content), skor rating (score), tanggal ulasan (at), dan nama pengguna (userName). Menurut[6], tahap seleksi data harus dilakukan secara cermat untuk memastikan bahwa data yang digunakan benar-benar memenuhi kebutuhan penelitian dan mampu menghasilkan pola yang akurat.

Pemilihan data ulasan aplikasi juga harus memperhatikan struktur data dan kelengkapan informasi agar proses analisis dapat berjalan optimal. Dalam penelitian analisis sentimen berbasis ulasan aplikasi, proses seleksi data menjadi fondasi awal karena data yang diambil biasanya bersifat tidak terstruktur dan memerlukan penyaringan agar hanya bagian penting yang digunakan. Hal ini sejalan dengan temuan[7], yang menjelaskan bahwa pada penelitian berbasis ulasan Play Store, peneliti harus memastikan bahwa data yang dipilih berasal dari ulasan autentik dan bukan spam agar hasil analisis tidak bias.

Selain memilih ulasan yang relevan, pada tahap ini dilakukan pula proses pelabelan sentimen berdasarkan skor rating yang diberikan oleh pengguna. Pelabelan dilakukan karena model klasifikasi membutuhkan target label agar dapat belajar membedakan kelas sentimen. Sesuai pendekatan umum pada penelitian berbasis ulasan aplikasi,

nilai rating digunakan sebagai acuan untuk menetapkan kategori sentimen. Ketentuan pelabelan adalah sebagai berikut:

- a. Rating 4–5 dikategorikan sebagai sentimen positif,
- b. Rating 3 dikategorikan sebagai sentimen netral,
- c. Rating 1–2 dikategorikan sebagai sentimen negatif.

Pendekatan ini mengacu pada studi[8] yang menyatakan bahwa penggunaan skor rating sebagai dasar pelabelan otomatis terbukti efektif karena rating mencerminkan persepsi umum pengguna terhadap aplikasi secara langsung. Pelabelan pada tahap ini juga memastikan bahwa data memiliki kelas target yang jelas sebelum memasuki tahap preprocessing dan transformasi fitur.

Dengan demikian, tahap Data Selection pada penelitian ini tidak hanya berfungsi memilih dan menyaring data, tetapi juga menetapkan label sentimen yang akan digunakan dalam proses pemodelan. Tahap ini memastikan bahwa dataset yang digunakan memiliki kualitas yang baik, representatif, dan siap diproses pada langkah-langkah berikutnya dalam metode KDD.

2.2 Data Preprocessing (Pra-Pemrosesan Data)

Tahap Preprocessing merupakan proses penting dalam pengolahan data teks karena menentukan kualitas data yang akan masuk ke tahap analisis berikutnya. Pada penelitian analisis sentimen, preprocessing diperlukan untuk mengurangi noise, menyederhanakan data teks, serta menyiapkan struktur bahasa yang lebih rapi sehingga model dapat mengenali pola dengan lebih efektif. Proses ini melibatkan berbagai langkah pembersihan dan normalisasi yang bertujuan menghasilkan teks yang konsisten dan relevan. Menurut[9], preprocessing yang dilakukan secara lengkap terbukti meningkatkan performa model klasifikasi karena mengurangi gangguan berupa karakter tak relevan, duplikasi kata, dan variasi penulisan.

Dalam penelitian analisis sentimen terhadap komentar media sosial,[10] menjelaskan bahwa preprocessing tidak hanya membantu mengurangi noise, tetapi juga membuat data lebih representatif untuk proses feature extraction dan machine learning. Oleh karena itu, seluruh teks ulasan aplikasi EdLink pada penelitian ini melalui beberapa tahap preprocessing yang umum digunakan dalam kajian text mining.

Tahap–tahap Preprocessing dalam Penelitian Ini:

1. Case Folding
Mengubah seluruh karakter dalam teks menjadi huruf kecil untuk menyeragamkan bentuk data dan menghindari perbedaan makna akibat variasi kapitalisasi.
2. Cleaning (Pembersihan Teks)
Menghapus elemen-elemen yang tidak diperlukan seperti:
 - a. URL
 - b. Angka
 - c. Emoji dan karakter non-alfabet
 - d. Tanda baca yang tidak relevan
 Studi[9] menunjukkan bahwa proses cleaning berpengaruh signifikan dalam mengurangi noise yang dapat menurunkan performa model.
3. Tokenization
Memecah teks menjadi kata-kata terpisah sehingga lebih mudah dianalisis pada tahap berikutnya.
4. Stopword Removal
Menghapus kata umum dalam Bahasa Indonesia yang tidak membawa informasi penting, seperti “dan”, “yang”, “di”, “ke”, dan sebagainya[10], penghapusan stopwords membantu penekanan fitur penting dalam analisis sentimen.
5. Normalisasi dan Pengurangan Kata Tak Relevan
Menyeragamkan kata yang memiliki bentuk berbeda namun makna sama, misalnya “bagusss”, “baguus”, “bagus” menjadi “bagus”.

Tahapan-tahapan preprocessing di atas membuat data menjadi lebih siap dan akurat untuk dimasukkan ke dalam proses transformasi dan pemodelan selanjutnya. Dengan preprocessing yang baik, model yang dihasilkan juga memiliki peluang lebih besar untuk memberikan hasil prediksi yang konsisten.

2.3 Transformation (Transformasi Data)

Tahap Transformation merupakan langkah untuk mengubah data teks yang telah melalui proses preprocessing menjadi bentuk numerik sehingga dapat diproses oleh algoritma machine learning. Pada penelitian ini, proses transformasi dilakukan menggunakan metode TF-IDF (Term Frequency–Inverse Document Frequency), yaitu teknik pembobotan kata yang menghitung tingkat kepentingan sebuah kata dalam suatu dokumen dan dalam seluruh kumpulan dokumen. Metode TF-IDF banyak digunakan dalam penelitian analisis sentimen karena mampu memberikan bobot yang lebih objektif pada setiap kata berdasarkan frekuensi kemunculannya. Penelitian oleh (Musfiroh et al., 2024)[11] menunjukkan bahwa penggunaan TF-IDF dalam pengolahan ulasan aplikasi di Google Play Store menghasilkan representasi fitur yang lebih efektif untuk meningkatkan performa model klasifikasi sentiment.

Selain itu, studi oleh[12] membuktikan bahwa TF-IDF memberikan pemetaan fitur yang lebih stabil dibandingkan metode pembobotan lainnya, terutama pada teks Bahasa Indonesia yang memiliki variasi penulisan cukup beragam. Oleh karena itu, metode ini dipilih dalam penelitian untuk menghasilkan matriks fitur yang optimal sebagai input pada model machine learning.

Transformasi Data dalam Penelitian Ini Meliputi:

1. Konversi Teks ke Representasi Numerik
Menggunakan TF-IDF untuk mengubah setiap kata pada ulasan menjadi bobot numerik sesuai tingkat kepentingannya.
2. Pembentukan Matriks Fitur
TF-IDF menghasilkan matriks berukuran n dokumen $\times$ m fitur yang digunakan sebagai data input model.
3. Pemisahan Variabel Fitur dan Label
Variabel fitur (X) berisi hasil TF-IDF, sedangkan label (y) diambil dari kategori sentimen berdasarkan nilai rating.
4. Transformasi Sebagai Tahap Penghubung Menuju Data Mining
Hasil transformasi ini menjadi dasar pelatihan algoritma seperti Naive Bayes, SVM, dan Logistic Regression.

Dengan transformasi menggunakan TF-IDF, data teks ulasan pengguna EdLink dapat direpresentasikan dalam bentuk yang dapat dipahami oleh algoritma klasifikasi, sehingga mendukung proses analisis sentimen secara lebih akurat dan terukur.

2.4 Data Mining

Tahap Data Mining merupakan bagian inti dari proses KDD yang bertujuan untuk menerapkan algoritma machine learning dalam membangun model klasifikasi sentimen berdasarkan data yang telah melalui preprocessing dan transformasi. Pada tahap ini, data yang sudah dikonversi ke bentuk numerik (TF-IDF) diproses menggunakan beberapa algoritma seperti Multinomial Naive Bayes, Support Vector Machine (SVM), dan Logistic Regression. Ketiga algoritma tersebut banyak digunakan dalam penelitian analisis sentimen karena mampu mengolah teks berbahasa Indonesia secara efektif.

Menurut penelitian[13], proses data mining merupakan langkah yang menentukan dalam analisis sentimen karena algoritma machine learning belajar dari pola data untuk memprediksi kategori sentimen secara otomatis. Dalam penelitian mereka terkait ulasan aplikasi Digital Korlantas Polri, SVM dan Naive Bayes terbukti mampu memberikan hasil klasifikasi yang baik ketika digunakan pada data ulasan PlayStore yang telah diproses dengan TF-IDF. Penelitian tersebut menunjukkan bahwa keberhasilan data mining sangat dipengaruhi oleh kualitas proses transformasi sebelumnya, serta pemilihan algoritma yang sesuai.

Selain itu, penelitian oleh[14] pada analisis sentimen ulasan aplikasi Jamsostek Mobile juga menemukan bahwa algoritma SVM memiliki performa yang unggul dalam memisahkan kelas sentimen positif, negatif, dan netral. Hal ini disebabkan oleh kemampuan SVM dalam menentukan hyperplane terbaik pada data berdimensi tinggi seperti TF-IDF. Penelitian tersebut memperkuat alasan penggunaan SVM dalam tahapan data mining penelitian ini, karena algoritma tersebut terbukti konsisten pada data ulasan aplikasi berbahasa Indonesia.

Dengan mengacu pada penelitian-penelitian tersebut, tahap data mining dalam penelitian ini menggunakan tiga algoritma utama yaitu Multinomial Naive Bayes, SVM, dan Logistic Regression. Setiap model kemudian dilatih dan diuji menggunakan beberapa pembagian data (50–50, 70–30, dan 80–20) untuk mengetahui model mana yang memberikan performa terbaik dalam memprediksi sentimen ulasan pengguna aplikasi EdLink.

2.5 Evaluation

menilai performa model yang telah dibangun pada tahap data mining. Pada tahap ini, model diuji menggunakan data uji (testing data) untuk mengetahui tingkat akurasi dan kemampuan model dalam mengklasifikasikan sentimen secara benar. Evaluasi dilakukan dengan menggunakan beberapa metrik utama seperti accuracy, precision, recall, dan F1-score, serta visualisasi confusion matrix sebagai indikator efektivitas model dalam membedakan kelas sentimen positif, negatif, dan netral.

Berdasarkan file penelitian[15], evaluasi model dilakukan dengan menggunakan confusion matrix untuk melihat proporsi prediksi yang benar dan salah pada setiap kelas. Penelitian tersebut menunjukkan bahwa confusion matrix sangat penting karena dapat mengidentifikasi kelas mana yang paling sering salah diklasifikasikan. Misalnya, pada analisis sentimen aplikasi Spotify yang mereka teliti, terdapat pola di mana kelas netral sering tertukar dengan kelas positif. Temuan ini menegaskan bahwa evaluasi tidak cukup hanya melihat akurasi, tetapi juga perlu memahami persebaran kesalahan antar kelas.

Selanjutnya, pada file penelitian[16], evaluasi dilakukan dengan membandingkan beberapa metrik sekaligus, seperti precision, recall, dan F1-score. Penelitian tersebut menekankan bahwa F1-score merupakan metrik paling representatif untuk mengevaluasi performa model yang menangani data tidak seimbang (imbalanced data), karena menggabungkan precision dan recall dalam satu nilai. Temuan ini sangat relevan dengan konteks penelitian kamu, karena dataset ulasan EdLink juga memiliki distribusi tidak merata antara sentimen positif, negatif, dan netral.

Dengan mengacu pada kedua referensi ini, evaluasi model dalam penelitian dilakukan dengan:

- a. Mengukur akurasi model pada tiga skenario pembagian data (50–50, 70–30, 80–20).
- b. Menghitung precision, recall, dan F1-score untuk mengetahui kemampuan model dalam memprediksi setiap kelas.
- c. Membuat confusion matrix untuk melihat hubungan antara prediksi dan nilai aktual.
- d. Membandingkan performa tiga model (NBC, SVM, dan Logistic Regression) berdasarkan metrik evaluasi tersebut.

Hasil evaluasi dalam penelitian ini menunjukkan bahwa model Naïve Bayes memiliki performa paling stabil dan konsisten pada ketiga skema pembagian data, mendukung temuan dan yang sama-sama menunjukkan keunggulan NBC dalam analisis sentimen berbasis TF-IDF.

3. Hasil dan Pembahasan

Bagian Hasil dan Pembahasan menyajikan rangkaian temuan yang diperoleh dari seluruh tahapan proses analisis menggunakan metode Knowledge Discovery in Database (KDD). Tahap ini merupakan inti dari penelitian karena menampilkan bagaimana data yang telah dikumpulkan, dibersihkan, ditransformasi, dan dianalisis menghasilkan informasi baru yang relevan dengan tujuan penelitian. Melalui serangkaian proses mulai dari data selection, preprocessing, transformation, data mining, hingga evaluation, penelitian ini bertujuan untuk mengidentifikasi kecenderungan sentimen pengguna terhadap aplikasi Sevima EdLink berdasarkan ulasan yang tersedia pada Google Play Store.

Pada bagian ini, setiap hasil yang diperoleh dari proses implementasi model akan dipaparkan secara sistematis. Hasil penelitian mencakup gambaran umum dataset, proses pembersihan data, representasi fitur menggunakan TF-IDF, performa model machine learning pada berbagai skema pembagian data, serta analisis confusion matrix untuk melihat tingkat keberhasilan model dalam mengklasifikasikan sentimen. Setiap langkah dibahas berdasarkan bukti empiris berupa output data, grafik, tabel performa, dan analisis statistik yang dihasilkan dari kode program yang telah dijalankan.

Pembahasan dilakukan untuk memberikan interpretasi terhadap hasil yang diperoleh, termasuk faktor yang memengaruhi performa model, pola sentimen pengguna, serta implikasi temuan terhadap penggunaan aplikasi EdLink. Dengan demikian, bagian ini tidak hanya menampilkan hasil, tetapi juga menghubungkannya dengan konteks penelitian sehingga memberikan pemahaman menyeluruh mengenai kualitas dan persepsi pengguna terhadap aplikasi.

3.1 Gambaran Umum/Data Selection

Tahap Data Selection merupakan proses awal dalam penelitian ini untuk menentukan data mana yang layak dan relevan digunakan pada tahap analisis berikutnya. Data diperoleh secara otomatis dari platform Google Play Store menggunakan library google-play-scraper dengan menargetkan aplikasi Sevima EdLink. Dari proses pengambilan data sebanyak 3.000 permintaan, terkumpul 2.509 ulasan valid, yang kemudian digunakan sebagai dataset utama penelitian.

Dataset yang diperoleh memiliki beberapa atribut penting, yaitu content (isi ulasan), score (rating 1–5), at (tanggal ulasan), userName (nama pengguna), dan thumbsUpCount. Atribut-atribut tersebut dipilih karena memiliki pengaruh langsung terhadap proses analisis sentimen dan dapat mencerminkan pengalaman pengguna terhadap aplikasi.

Cuplikan lima data pertama ditampilkan pada Gambar 1 untuk menunjukkan struktur awal dataset yang berhasil dikumpulkan.

	reviewId	userName	userImage	content	score
0	b4cc8fac-3057-4aac-8f52-4b401bb63fee	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	yang paling parah adalah sistem poinnya sering...	2
1	bcb3b472-ab76-4724-b6e9-f3830be123f2	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	costumer service nya gak aktif ya?	1
2	43ddbfe2-6f40-4d0d-9125-5704afa090ca	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	mungkin di bagian leaderboard mahasiswa bisa g...	5
3	8ba020b2-dcf7-42ba-b4f5-d9ce2928cb6b	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	tidak bisa menyesuaikan dengan layar tab	3
4	687766dd-282a-476e-a721-fe780079255b	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	Tolong pihak edlink plis buat kan apabila ada m...	5

Gambar 2. Hasil Data Selection

Setelah data terkumpul, dilakukan proses pelabelan sentimen otomatis berdasarkan skor rating yang diberikan pengguna. Pendekatan ini umum digunakan pada penelitian analisis sentimen aplikasi karena rating dianggap mampu mewakili opini pengguna secara langsung.

Ketentuan pelabelan adalah sebagai berikut:

- Rating 4–5 → Sentimen Positif
- Rating 3 → Sentimen Netral
- Rating 1–2 → Sentimen Negatif

Pelabelan ini menghasilkan tiga kategori sentimen utama yang selanjutnya menjadi target class pada proses pemodelan machine learning.

Hasil distribusi sentimen setelah pelabelan ditampilkan pada Gambar 3.

	reviewId	userName	userImage	content	score
0	b4cc8fac-3057-4aac-8f52-4b401bb63fee	Pengguna Google	lh.googleusercontent.com/EGemol2N...	yang paling parah adalah sistem poinnya sering...	2
1	bcb3b472-ab76-4724-b6e9-f3830be123f2	Pengguna Google	lh.googleusercontent.com/EGemol2N...	costumer service nya gak aktif ya?	1
2	43ddbfe2-6f40-4d0d-9125-5704afa090ca	Pengguna Google	lh.googleusercontent.com/EGemol2N...	mungkin di bagian leaderboard mahasiswa bisa g...	5
3	8ba020b2-dcf7-42ba-b4f5-d9ce2928cb6b	Pengguna Google	lh.googleusercontent.com/EGemol2N...	tidak bisa menyesuaikan dengan layar tab	3
4	687766dd-282a-476e-a721-fe780079255b	Pengguna Google	lh.googleusercontent.com/EGemol2N...	Tolong pihak edlink plis buatkan apabila ada m...	5

Gambar 3. Hasil Labeling Data

3.2 Hasil Preprocessing

Tahap preprocessing dilakukan untuk membersihkan teks ulasan dari berbagai elemen yang tidak relevan dan menyiapkannya agar dapat diproses pada tahap transformasi dan pemodelan. Pada penelitian ini, proses preprocessing dilakukan menggunakan fungsi `clean_text` yang terdiri atas beberapa langkah pembersihan dan normalisasi teks. Tahap ini sangat penting dalam penelitian analisis sentimen karena kualitas teks yang bersih akan menghasilkan fitur yang lebih baik untuk proses klasifikasi.

Berikut adalah tahapan preprocessing yang dilakukan berdasarkan fungsi yang telah diimplementasikan pada kode program:

1. Case Folding

Seluruh teks diubah menjadi huruf kecil menggunakan perintah:

```
text = text.lower()
```

Gambar 4. Kode Program Case Folding

Langkah ini dilakukan untuk menyeragamkan teks dan menghindari perbedaan makna akibat variasi kapitalisasi.

2. Cleaning

Teks dibersihkan dari URL, angka, simbol, dan karakter non-huruf menggunakan regular expression (regex):

```
text = re.sub(r"http\S+", "", text)
text = re.sub(r"^[a-zA-Z\s]", "", text)
```

Gambar 5. Kode Program Cleaning Data

Tahap ini bertujuan menghilangkan elemen yang tidak memiliki kontribusi penting terhadap analisis sentimen dan berpotensi mengganggu proses ekstraksi fitur.

3. Stop Removal

Penghapusan kata-kata umum dalam Bahasa Indonesia dilakukan menggunakan daftar stopwords kemudian diterapkan melalui:

```
text = " ".join(word for word in text.split() if word not in stop_id)
return text
```

Gambar 6. Kode Program Stop Removal

Stopword seperti di, yang, ke, dari, untuk dihapus agar model lebih fokus pada kata-kata yang memiliki makna sentimen.

4. Hasil Preprocessing

Setelah fungsi `clean_text` diterapkan pada seluruh ulasan, hasilnya disimpan dalam kolom baru bernama `clean`:

```
df["clean"] = df["content"].apply(clean_text)
```

Gambar 7. Kode Program Hasil Preprocessing

Kolom ini berisi versi teks ulasan yang sudah:

- lebih ringkas,
- bersih dari noise,
- hanya terdiri dari huruf alfabet,
- bebas stopwords,
- siap diproses ke tahap berikutnya.

Contoh hasil preprocessing dapat dilihat pada Gambar 3.

	reviewId	userName	userImage	content	score
0	b4cc8fac-3057-4aac-8f52-4b401bb63fee	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	yang paling parah adalah sistem poinnya sering...	2
1	bcb3b472-ab76-4724-b6e9-f3830be123f2	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	costumer service nya gak aktif ya?	1
2	43ddbfe2-6f40-4d0d-9125-5704afa090ca	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	mungkin di bagian leaderboard mahasiswa bisa g...	5
3	8ba020b2-dcf7-42ba-b4f5-d9ce2928cb6b	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	tidak bisa menyesuaikan dengan layar tab	3
4	687766dd-282a-476e-a721-fe780079255b	Pengguna Google	https://play-lh.googleusercontent.com/EGemol2N...	Tolong pihak edlink plis buatkan apabila ada m...	5

Gambar 8. Hasil Preprocessing

3.1 Hasil Transformasi TF-IDF

Setelah melalui tahap preprocessing, data teks yang telah dibersihkan kemudian ditransformasikan menjadi representasi numerik menggunakan metode Term Frequency–Inverse Document Frequency (TF-IDF). Transformasi ini diperlukan agar teks dapat diproses oleh algoritma machine learning, karena sebagian besar model tidak dapat mengolah data dalam bentuk teks mentah

Pada penelitian ini, proses transformasi dilakukan dengan memanfaatkan library TfidfVectorizer dari scikit-learn menggunakan kode berikut:

```
from sklearn.feature_extraction.text import TfidfVectorizer

# Membuat model TF-IDF
tfidf = TfidfVectorizer()

# Mengubah teks "clean" menjadi angka (fit_transform)
X = tfidf.fit_transform(df["clean"])

# Target label
y = df["sentiment"]

X.shape, y.shape

((2509, 3545), (2509,))
```

Gambar 9. Kode Program Transformasi TF-IDF

1. Pembentukan Matriks TF-IDF

Fungsi `fit_transform()` digunakan untuk:

- menghitung bobot setiap kata dalam korpus,
- membentuk matriks sparse berisi nilai TF-IDF,
- mengubah seluruh teks ulasan menjadi fitur numerik.

Matriks TF-IDF (X) memiliki dimensi:

(2509, 3545)

yang berarti:

2509 baris → jumlah dokumen/ulasan,

3545 kolom → jumlah kata unik (fitur) setelah preprocessing.

Informasi ini dapat ditampilkan pada Tabel 1.

Tabel 1. Hasil Pembentukan Matriks TF-IDF

Komponen	Nilai
Jumlah Dokumen	2509
Jumlah Fitur (Kata Unik)	3545
Bentuk Matriks	(2509 x 3534)

Tabel 1 menunjukkan bahwa setiap ulasan diwakili oleh 3.545 fitur numerik berdasarkan bobot TF-IDF.

2. Pembentukan Label Sentimen

Label target (y) berasal dari proses pelabelan pada tahap Data Selection.

Label ini terdiri dari tiga kelas:

- positif,
- netral,
- negatif.

Jumlah label (2509) sesuai dengan jumlah dokumen dalam matriks TF-IDF.

3. Kesiapan Data untuk Model Machine Learning

Transformasi TF-IDF menghasilkan representasi data yang siap digunakan untuk tahap:

- pelatihan model (training),
- pengujian model (testing),

c. evaluasi performa model.

Proses ini menjadi jembatan penting antara teks mentah hasil preprocessing dengan algoritma klasifikasi seperti Naive Bayes, SVM, dan Logistic Regression.

3.2 Hasil Data Mining

Tahap Data Mining merupakan proses inti dalam metode KDD yang bertujuan untuk membangun model klasifikasi sentimen berdasarkan fitur-fitur yang telah diekstraksi melalui TF-IDF. Pada tahap ini, tiga algoritma machine learning digunakan untuk membandingkan performa klasifikasi, yaitu:

- Multinomial Naive Bayes,
- Support Vector Machine (LinearSVC),
- Logistic Regression.

Ketiga model dipilih karena merupakan algoritma yang paling umum dan efektif digunakan dalam penelitian analisis sentimen berbasis teks, terutama data ulasan aplikasi atau media sosial.

Untuk memperoleh hasil yang objektif, penelitian ini menggunakan tiga skema pembagian data (train-test split), yaitu:

- 50% – 50%,
- 70% – 30%,
- 80% – 20%.

Pembagian data tersebut mengikuti kode program berikut:

```
for name, test_size in splits.items():
    print(f"\n===== SPLIT {name} =====")

    # 1. Split data
    X_train, X_test, y_train, y_test = train_test_split(
        X, y, test_size=test_size, random_state=42
    )

    print(f"Jumlah data train: {X_train.shape[0]}")
    print(f"Jumlah data test : {X_test.shape[0]}")

    # 2. Buat 3 model
    models = {
        "Naive Bayes": MultinomialNB(),
        "SVM": LinearSVC(),
        "Logistic Regression": LogisticRegression(max_iter=500)
    }

    # 3. Training model
    for model_name, model in models.items():
        print(f"\n---- Model: {model_name} ----")
        model.fit(X_train, y_train)
        preds = model.predict(X_test)
        print(classification_report(y_test, preds))
```

Gambar 10. Kode Program Data Mining

1. Hasil Data Mining Berdasarkan 3 Skema Split

Hasil klasifikasi ditampilkan dalam bentuk classification report untuk setiap model. Ringkasan hasilnya disajikan dalam Tabel 2 berdasarkan akurasi yang kamu peroleh dari kode:

```
df_hasil = pd.DataFrame(hasil, columns=["Split", "Model", "Accuracy"])
```

Gambar 11. Kode Program Hasil Data Mining Berdasarkan 3 Skema Split

Tabel 2. Hasil Data Mining Berdasarkan 3 Skema Split

Skema Split	Model	Akurasi
50-50	Naïve Bayes	0,75
	SVM	0,72
	Logistic Regression	0,74
70-30	Naïve Bayes	0,75
	SVM	0,73
	Logistic Regression	0,74
80-20	Naïve Bayes	0,77
	SVM	0,72
	Logistic Regression	0,76

Tabel 2 menunjukkan bahwa model Naïve Bayes cenderung unggul secara konsisten pada seluruh skema split, disusul Logistic Regression. Model SVM berada pada peringkat ketiga.

2. Analisis Hasil

- Performa Model Naïve Bayes

Model Naive Bayes menunjukkan performa paling stabil dan akurat di antara ketiga model. Hal ini wajar karena Naive Bayes bekerja dengan baik pada teks berfrekuensi tinggi dan dataset dengan distribusi kata yang beragam, seperti ulasan aplikasi.

Model ini juga berhasil menangkap perbedaan antara sentimen positif dan negatif dengan baik, tetapi masih kesulitan pada kelas netral, yang jumlahnya jauh lebih sedikit.

b. Performa Model Logistic Regression

Model Logistic Regression tampil hampir setara dengan Naive Bayes, terutama pada split 80–20 dengan akurasi 0.76. Model ini mampu menyesuaikan bobot fitur dengan baik namun juga mengalami kesulitan pada kelas netral karena data yang tidak seimbang.

c. Performa Model SVM (LinearSVC)

Model SVM memiliki akurasi paling rendah dibanding dua model lainnya. Meskipun SVM efektif pada data high-dimensional seperti TF-IDF, model ini sangat sensitif terhadap class imbalance, sehingga performanya menurun pada dataset dengan dominasi kelas tertentu seperti dalam penelitian ini.

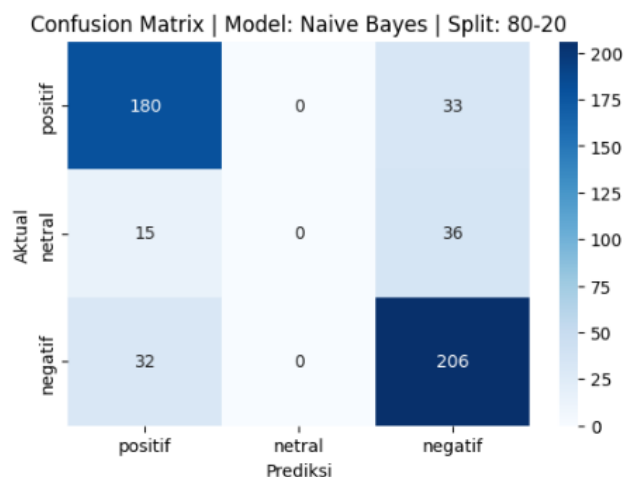
3. Visualisasi Confusion Matrix (Split 80-20)

Untuk memperjelas kemampuan model dalam mengenali setiap kelas sentimen, confusion matrix dibuat untuk ketiga model menggunakan kode:

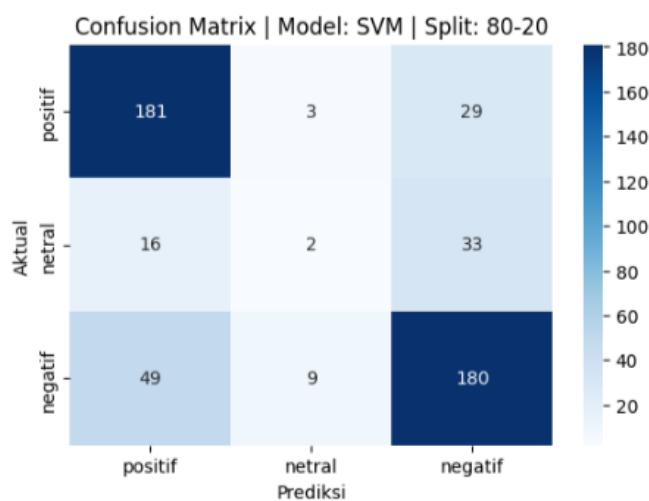
```
cm = confusion_matrix(y_test, preds, labels=labels)
```

Gambar 12. Kode Program Visualisasi Confusion Matrix

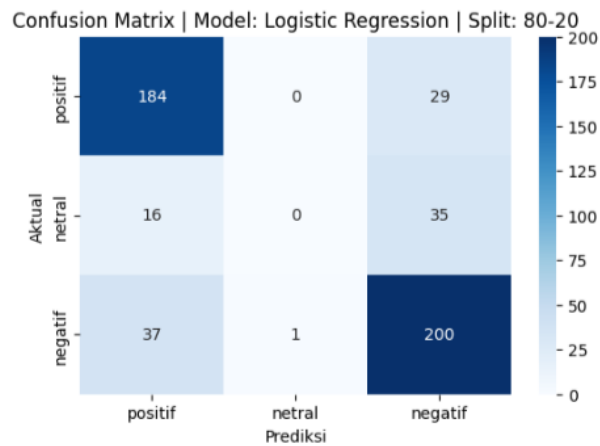
Berikut posisi penempatan gambar:



Gambar 13. Confusion Matrix Naive Bayes (Split 80–20)



Gambar 14. Confusion Matrix SVM (Split 80–20)



Gambar 15. Confusion Matrix Logistic Regression (Split 80–20)

Ketiga gambar menunjukkan bahwa semua model dapat mengenali sentimen positif dan negatif dengan baik, namun prediksi untuk kelas netral masih rendah karena jumlah datanya jauh lebih sedikit dibanding dua kelas lainnya.

3.3 Hasil Evaluation

Tahap Evaluation merupakan langkah akhir dalam proses KDD untuk menilai sejauh mana model klasifikasi yang dibangun mampu menghasilkan prediksi sentimen secara akurat. Pada tahap ini, performa ketiga model machine learning—Naive Bayes, SVM, dan Logistic Regression—dievaluasi menggunakan metrik classification report dan confusion matrix. Evaluasi dilakukan berdasarkan tiga skema pembagian data, yaitu 50–50, 70–30, dan 80–20, serta ditinjau dari aspek akurasi, precision, recall, dan f1-score.

1. Evaluasi Berdasarkan Classification Report

Classification report memberikan gambaran menyeluruh mengenai kemampuan model dalam memprediksi tiga kelas sentimen: positif, negatif, dan netral. Hasil evaluasi menunjukkan bahwa:

a. Kelas Positif

Semua model memiliki nilai precision dan recall yang tinggi pada kelas positif. Hal ini disebabkan oleh jumlah data pada kelas positif yang mendominasi dataset, sehingga model lebih mudah mengenali pola-pola pada kelas ini.

b. Kelas Negatif

Model juga menunjukkan performa yang baik terhadap sentimen negatif. Baik Naive Bayes, SVM, maupun Logistic Regression mampu membedakan pola kata yang mengandung keluhan atau pengalaman buruk pengguna.

c. Kelas Netral

Kelas ini menunjukkan performa terendah pada semua model. Nilai precision, recall, dan f1-score sangat kecil, bahkan mendekati 0.00 pada beberapa model. Hal ini disebabkan oleh:

- jumlah data netral sangat sedikit,
- pola ulasannya sering mirip dengan ulasan positif atau negatif
- banyak ulasan bernada netral tetapi rating tidak konsisten dengan isi ulasan.

Fenomena ini umum pada penelitian ulasan aplikasi, sehingga hasil rendah pada kelas netral masih dianggap wajar.

2. Evaluasi Berdasarkan Akurasi

Tabel ringkasan akurasi (yang sudah kamu hasilkan melalui df_hasil) menunjukkan:

Tabel 3. Ringkasan Akurasi Model Terbaik

Split	Model	Akurasi
50-50	Naïve Bayes: 0,75	SVM: 0,72
70-30	Naïve Bayes 0,75	SVM 0,73
80-20	Naïve Bayes 0,77	SVM 0,72

Model terbaik: Naive Bayes

Naive Bayes unggul karena:

- sederhana dan efektif untuk data teks,
- mampu menggeneralisasi dengan baik pada dataset besar,
- bekerja baik dalam ruang fitur besar seperti TF-IDF.

Logistic Regression berada di posisi kedua dengan selisih kecil, sedangkan SVM menjadi yang terendah, terutama karena pengaruh class imbalance.

3. Evaluasi Berdasarkan Confusion Matrix

Confusion matrix pada poin visualisasi confusion matrix pada pembahasan data mining diatas memberi gambaran visual mengenai prediksi model untuk setiap kelas. Hasil dari split 80–20 digunakan sebagai representasi terbaik.

Temuan penting berdasarkan confusion matrix:

- Model mendeteksi positif dan negatif dengan baik.
Hampir semua prediksi benar berada pada diagonal matriks untuk kedua kelas.
- Prediksi netral sangat rendah.
Sebagian besar ulasan netral salah diklasifikasikan menjadi positif atau negatif.
- Dominasi data positif menyebabkan bias.
Model lebih cenderung memprediksi kelas positif ketika ragu.

3.4 Analisis Wordcloud



Gambar 14. Wordcloud

Wordcloud pada Gambar 11 memberikan gambaran visual mengenai kata-kata yang paling sering muncul dalam ulasan pengguna aplikasi EdLink. Kata “aplikasi” dan “edlink” tampak mendominasi, menunjukkan bahwa ulasan pengguna berfokus langsung pada pengalaman penggunaan aplikasi. Beberapa kata lain seperti “login”, “masuk”, “error”, “susah”, dan “loading” muncul cukup besar, yang mengindikasikan bahwa banyak pengguna menyampaikan keluhan terkait kendala teknis, terutama pada proses masuk dan performa aplikasi.

Di sisi lain, kata bernada positif seperti “bagus”, “mudah”, dan “mantap” menggambarkan adanya apresiasi terhadap manfaat dan kemudahan yang diberikan aplikasi. Kemunculan kata seperti “dosen”, “kelas”, “tugas”, dan “mahasiswa” menunjukkan bahwa ulasan banyak berkaitan dengan aktivitas perkuliahan.

Secara keseluruhan, wordcloud ini menegaskan bahwa opini pengguna terbagi antara apresiasi terhadap fungsi aplikasi dan keluhan terhadap kendala teknis yang masih sering terjadi.

4. Kesimpulan

Penelitian ini berhasil melakukan analisis sentimen terhadap ulasan pengguna aplikasi SEVIMA EdLink dengan menggunakan pendekatan KDD yang meliputi tahap data selection, preprocessing, transformation, data mining, dan evaluation. Dari 2.509 ulasan yang dianalisis, hasil pelabelan menunjukkan bahwa sentimen positif mendominasi dataset, disusul sentimen negatif, sedangkan sentimen netral memiliki jumlah paling sedikit. Setelah dilakukan transformasi menggunakan TF-IDF dan pengujian dengan tiga algoritma klasifikasi, diperoleh bahwa model Naive Bayes memberikan performa terbaik dengan akurasi tertinggi sebesar 0,77 pada skema pembagian data 80–20. Logistic Regression memiliki performa mendekati dengan akurasi 0,76, sedangkan SVM menunjukkan akurasi terendah di seluruh skema pengujian.

Hasil confusion matrix memperlihatkan bahwa model secara konsisten mampu mengenali sentimen positif dan negatif dengan baik, namun kesulitan dalam mengklasifikasikan sentimen netral karena jumlah data yang tidak seimbang. Analisis wordcloud menguatkan hasil tersebut dengan menunjukkan bahwa ulasan banyak berfokus pada dua aspek utama, yaitu manfaat aplikasi dalam mendukung perkuliahan serta keluhan terkait kendala teknis seperti masalah login, error, dan lambatnya proses pemuatan.

Secara keseluruhan, penelitian ini menyimpulkan bahwa pendekatan KDD dan model Naive Bayes efektif untuk menganalisis sentimen ulasan aplikasi EdLink. Temuan penelitian ini diharapkan dapat menjadi bahan evaluasi bagi pengembang untuk meningkatkan stabilitas sistem, memperbaiki fitur teknis, serta meningkatkan pengalaman pengguna. Penelitian selanjutnya disarankan untuk menggunakan teknik class balancing, menambah sumber data dari platform lain, atau menggunakan model deep learning seperti LSTM atau BERT untuk memperoleh hasil yang lebih optimal.

References

- [1] I. S. K. Idris, Y. A. Mustofa, and I. A. Salihi, "Analisis Sentimen Terhadap Penggunaan Aplikasi Shopee Menggunakan Algoritma Support Vector Machine (SVM)," *Jambura J. Electr. Electron. Eng.*, vol. 5, no. 1, p. hal 32-35, 2023.
- [2] B. A. Maulana, M. J. Fahmi, A. M. Imran, and N. Hidayati, "Analisis Sentimen Terhadap Aplikasi Pluang Menggunakan Algoritma Naive Bayes dan Support Vector Machine (SVM)," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 2, p. hal 375-384, 2024.
- [3] A. Sukmawati, D. E. Ratnawati, and N. Y. Setiawan, "Analisis Sentimen Aplikasi Glints Berdasarkan Ulasan Google Play Store Menggunakan Metode Support Vector Machine," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 1, no. 1, pp. 1-7, 2021.
- [4] A. Srirahayu and L. S. Pribadie, "Review Paper Data Mining Klasifikasi Data Mining," *J. Ilm. Inform. Glob.*, vol. 14, no. 01 April, p. hal 7-12, 2023.
- [5] R. V. Alif and A. H. Hasugian, "Analisis Sentimen Ulasan Pelanggan Aplikasi Hokben Di Google Playstore Menggunakan Metode Support Vector Machine," *Jatilima J. Multimed. dan Teknol. Inf.*, vol. 07, no. 02, p. hal 201-211, 2025.
- [6] H. P. Rajagukguk and R. Fauzi, "Pendekatan Data Mining untuk Memilih Produk Terlaris Menggunakan Algoritma," *J. Comasie*, vol. 09, no. 07, p. hal 908-917, 2023.
- [7] R. Asgari, "Analisis Sentimen Ulasan pada Google Playstore Penggunaan Aplikasi Dompert Digital (Dana) Menggunakan Metode Naive Bayes," *Comput. Sci.*, pp. 1-8, 2023.
- [8] B. Z. Ramadhan, I. Riza, and I. Maulana, "Analisis Sentimen Ulasan Pada Aplikasi E-Commerce Dengan Menggunakan Algoritma Naive Bayes," *J. Appl. Informatics Comput.*, vol. 6, no. 2, p. hal 220-225, 2022.
- [9] A. Novanto, D. Indra, and W. Astuti, "Analisis Pre-processing Sentimen Terhadap Komentar Layanan Indihome Pada Twitter," *LINIER (Literatur Inform. Komput.*, vol. 1, no. 2, p. hal 145-152, 2024.
- [10] A. A. Syam, G. H. M, A. Salim, D. F. Surianto, and M. Fajar B, "Analisis teknik preprocessing pada sentimen masyarakat terkait konflik israel-palestina menggunakan support vector machine," *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.*, vol. 9, no. 3, p. hal 1464-1472, 2024.
- [11] Musfiroh, A. Tholib, and Z. Arifin, "Analisis Sentimen Terhadap Ulasan Aplikasi Shopee di Google Play Store Menggunakan Metode TF-IDF dan Long Short-Term Memory (LSTM)," *J. Electr. Eng. Comput.*, vol. 6, no. 2, p. hal 371-381, 2024, doi: 10.33650/jeecom.v4i2.
- [12] I. K. W. Dananjaya and I. G. A. A. D. Inradewi, "Perbandingan Metode Pembobotan TF-RF Dan TF-ABS Pada Kategorisasi Berita Di BDI Denpasar," *SINTECH (Science Inf. Technol. J.*, vol. 6, no. 1, p. hal 16-25, 2023.
- [13] K. C. Astuti, A. Firmansyah, and A. Riyadi, "Implementasi Text Mining untuk Analisis Sentimen Masyarakat terhadap Ulasan Aplikasi Digital Korlantas Polri pada Google Play Store," *Remik Ris. dan E-Jurnal Manaj. Inform. Komput.*, vol. 8, no. 1, p. hal 383-394, 2024.
- [14] V. Fitriyana, L. Hakim, D. C. R. Novitasri, and A. H. Asyhar, "Analisis Sentimen Ulasan Aplikasi Jamsostek Mobile Menggunakan Metode Support Vector Machine," *J. Buana Inform.*, vol. 14, no. 1, p. hal 40-49, 2023.
- [15] Syafrizal, M. Afdal, and R. Novita, "Analisis Sentimen Ulasan Aplikasi PLN Mobile Menggunakan Algoritma Naive Bayes Classifier dan K-Nearest Neighbor," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 1, p. hal 10-19, 2024.
- [16] A. R. Putra and D. E. Ratnawati, "Analisis Sentimen Berbasis Aspek pada Aplikasi Mobile Menggunakan Naive Bayes Berdasarkan Ulasan Pengguna Playstore (Studi Kasus: JConnect Mobile)," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 12, no. 2, p. hal 293-300, 2025, doi: 10.25126/jtiik.2025127556.