

Analisis Sentimen Evaluasi Pembelajaran Udemy Menggunakan *Support Vector Machine* Untuk Peningkatan Kursus

DOI: <http://dx.doi.org/10.35889/jutisi.v15i4.3968>

Creative Commons License 4.0 (CC BY – NC)

Irwansyah^{1*}, Ellya Helmut²

Sistem Informasi, ISB Atma Luhur, Pangkalpinang, Indonesia

*e-mail Corresponding Author: 2422520018@mahasiswa.atmaluhur.ac.id

Abstract

Learning evaluation plays an important role in improving course quality, yet user reviews are generally unstructured and difficult to analyze manually. This study aims to analyze user sentiment toward Udemy course evaluations using the Support Vector Machine (SVM) algorithm. The dataset consisted of 50 balanced English-language reviews (26 positive and 24 negative) selected from approximately 500 scraped reviews, with sentiment labels determined based on the rating categories provided by the source website. The research process included text preprocessing, TF-IDF feature weighting, SVM classification using a linear kernel, and 10-fold cross-validation. Model performance was evaluated using accuracy, precision, recall, and F1-score. The proposed model achieved an accuracy of 82% and an average F1-score of 0.82, indicating that the combination of TF-IDF and SVM effectively classifies user sentiment. These findings can assist course providers in evaluating user feedback and improving learning quality.

Keywords: Sentiment Analysis; Support Vector Machine; TF-IDF; Business Intelligence; Learning Evaluation

Abstrak

Evaluasi pembelajaran berperan penting dalam meningkatkan kualitas kursus, namun ulasan pengguna umumnya berupa data teks tidak terstruktur yang sulit dianalisis secara manual. Penelitian ini bertujuan menganalisis sentimen pengguna terhadap evaluasi pembelajaran pada platform Udemy menggunakan *algoritma Support Vector Machine* (SVM). Dataset terdiri atas 50 ulasan berbahasa Inggris yang seimbang (26 positif dan 24 negatif) dari sekitar 500 hasil *scraping*, dengan label sentimen ditentukan berdasarkan kategori rating pada situs sumber. Tahapan penelitian meliputi preprocessing teks, pembobotan TF-IDF, klasifikasi SVM berkernel linear, dan validasi menggunakan 10-fold *cross validation*. Kinerja model dievaluasi menggunakan *accuracy*, *precision*, *recall*, dan F1-score. Hasil penelitian menunjukkan *accuracy* sebesar 82% dan rata-rata F1-score sebesar 0,82, sehingga kombinasi TF-IDF dan SVM efektif mengklasifikasikan sentimen pengguna serta dapat mendukung peningkatan kualitas pembelajaran.

Kata Kunci: Analisis Sentimen; Support Vector Machine; TF-IDF; Business Intelligence; Evaluasi Pembelajaran

1. Pendahuluan

Perkembangan pesat platform pembelajaran daring telah mengubah cara masyarakat mengakses pendidikan, sehingga evaluasi terhadap kualitas pembelajaran menjadi aspek krusial yang perlu dikelola secara sistematis. Umpan balik pengguna, yang merefleksikan penilaian terhadap materi, metode pengajaran, hingga pengalaman belajar secara keseluruhan, merupakan sumber data yang sangat bernilai bagi penyelenggara kursus untuk terus meningkatkan mutu layanannya. Ketika dikelola dengan tepat, kumpulan opini tersebut dapat menjadi fondasi *Business Intelligence* (BI) yang mentransformasi data mentah menjadi wawasan strategis bagi pengambilan keputusan [1]. Oleh karena itu, kajian mengenai bagaimana ulasan pengguna dapat dianalisis secara otomatis dan sistematis menjadi topik yang penting untuk terus

dikembangkan, khususnya pada platform pembelajaran yang memiliki basis pengguna besar dan terus bertumbuh seperti Udemy.

Udemy tampil sebagai salah satu penyedia layanan edukasi daring terpopuler yang menawarkan aneka ragam kelas dari berbagai disiplin ilmu. Sejalan dengan lonjakan pendaftar dan variasi kelas yang ditawarkan, volume ulasan dari peserta turut mengalami peningkatan, memuat pandangan mengenai mutu bahan ajar, cara mengajar, keahlian pengajar, hingga impresi belajar secara keseluruhan. Namun demikian, mayoritas ulasan tersebut hadir dalam format teks mentah yang tidak terstruktur, sehingga peninjauannya secara manual menjadi sulit dan memakan waktu, terlebih ketika volume data sangat besar. Metode telaah konvensional juga rentan memicu bias subjektif pada hasil pengamatan. Akibatnya, wawasan esensial dari opini peserta jarang dimanfaatkan secara optimal sebagai landasan pengembangan sistem maupun penetapan kebijakan peningkatan mutu kursus [2]. Kondisi inilah yang melatarbelakangi kebutuhan terhadap metode otomatis yang mampu mentransformasi ulasan-ulasan acak tersebut menjadi rangkuman sentimen yang terstruktur dan mudah diinterpretasikan [3].

Berbagai riset telah berupaya menyelesaikan permasalahan analisis opini pengguna pada ranah pendidikan melalui pendekatan analisis sentimen. Kamil dkk. mengkaji pandangan konsumen terhadap aplikasi edukasi di Google Play menggunakan pemodelan topik serta analisis sentimen, dan membuktikan efektivitas pendekatan ini dalam melacak tren opini publik, meski kajian tersebut masih terbatas pada penilaian peranti lunak dan belum menyentuh mutu edukasi yang dirasakan peserta secara riil [4]. Sejalan dengan itu, Saninah dkk. menganalisis sentimen ulasan pengguna terhadap aplikasi pembelajaran bahasa Duolingo menggunakan algoritma *Naïve Bayes Classifier*, dan membuktikan bahwa metode tersebut mampu memilah ulasan pengguna secara akurat, namun kajian ini masih terbatas pada satu aplikasi pembelajaran bahasa tunggal dan belum menyentuh ranah kursus daring komersial lintas topik seperti Udemy [5]. Maulana dkk. menerapkan algoritma *Support Vector Machine* (SVM) untuk memetakan pandangan publik terhadap perguruan tinggi berdasarkan komentar YouTube, dan menegaskan keandalan SVM dalam mengkategorikan opini berbasis teks acak, namun sampel datanya berasal dari opini publik secara luas sehingga belum spesifik merepresentasikan pengalaman peserta kursus [6]. Rahayu dan Buditjahjanto menganalisis sentimen mahasiswa terhadap tugas perkuliahan menggunakan *K-Nearest Neighbors* (K-NN), dan menyimpulkan bahwa ekstraksi sentimen berguna untuk memantau reaksi peserta didik, meskipun K-NN kurang optimal dalam memproses data teks berdimensi kompleks dibandingkan SVM [7]. Kurnia dkk. mengevaluasi kepuasan mahasiswa terhadap aplikasi manajemen akademik berbasis web dengan analisis sentimen multikelas, dan memvalidasi bahwa klasifikasi sentimen mampu merepresentasikan tingkat kepuasan pengguna, tetapi fokusnya hanya menyasar sistem administratif dan belum menyentuh evaluasi mutu belajar-mengajar secara komprehensif [8]. Lubis dan Putri menganalisis sentimen mahasiswa terhadap penggunaan e-learning menggunakan SVM dan menyimpulkan algoritma ini mumpuni memisahkan opini ke kategori yang tepat, namun tinjauan tersebut hanya berpusat pada aspek penggunaan sistem daring dan belum memaksimalkan temuan sentimen sebagai pijakan peningkatan kualitas pengajaran [9]. Hanif dkk. membandingkan performa *Decision Tree*, *Random Forest*, dan SVM dalam mengklasifikasikan komentar mahasiswa, dan menunjukkan bahwa SVM memiliki performa kompetitif, namun penelitian tersebut lebih menitikberatkan pada perbandingan performa model dibandingkan pemanfaatan hasil analisis sentimen sebagai sumber informasi pendukung pengambilan keputusan [10]. Berdasarkan tinjauan tersebut, analisis sentimen telah banyak diterapkan pada bidang pendidikan untuk mengevaluasi layanan maupun sistem pembelajaran digital, tetapi penelitian yang secara khusus memanfaatkan ulasan pengguna kursus pada platform Udemy sebagai sumber informasi *Business Intelligence* guna mendukung peningkatan kualitas kursus masih relatif terbatas. Sebagian besar riset terdahulu juga berfokus pada perbandingan performa algoritma atau pada domain aplikasi/institusi pendidikan formal, dan belum menyasar ekosistem kursus daring komersial dengan basis ulasan lintas topik seperti Udemy. Celah (gap) inilah yang menjadi dasar dilakukannya penelitian ini.

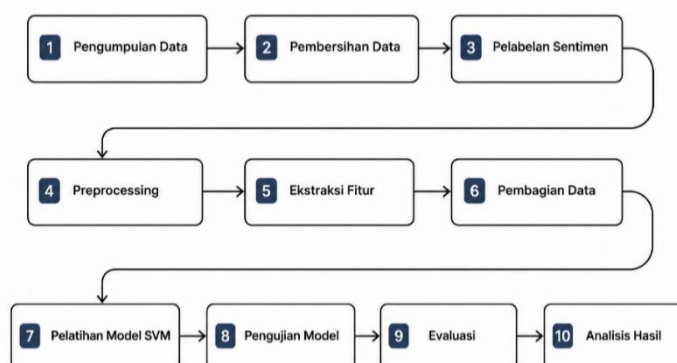
Untuk menjawab celah tersebut, penelitian ini mengusulkan penerapan algoritma *Support Vector Machine* (SVM) yang dikombinasikan dengan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengklasifikasikan sentimen ulasan pengguna kursus Udemy. Pemilihan SVM didasarkan pada kemampuannya menciptakan *hyperplane* optimal yang memisahkan kelas data secara maksimal, sehingga tangguh dalam menangani data teks berdimensi tinggi dan jarang (*sparse*), sebagaimana dibuktikan konsisten pada berbagai

riset klasifikasi sentimen [4], [9], [11]. Representasi TF-IDF dipilih karena mampu memberikan bobot lebih besar pada istilah yang secara khusus mencirikan suatu dokumen dibandingkan istilah yang umum ditemukan pada banyak dokumen, sehingga meningkatkan kualitas fitur yang menjadi masukan model klasifikasi [10], [12]. Kombinasi kedua pendekatan ini dipandang efektif menyelesaikan keterbatasan riset sebelumnya karena tidak hanya berfokus pada akurasi klasifikasi semata, melainkan juga diarahkan untuk menghasilkan wawasan *Business Intelligence* yang aplikatif bagi peningkatan mutu kursus [1]. Tujuan utama studi ini adalah menganalisis opini pengguna terkait efektivitas kursus UdeMy, memilah komentar ke dalam kelompok sentimen positif dan negatif melalui algoritma SVM, sekaligus memetakan elemen edukasi yang kerap memicu reaksi memuaskan maupun sebaliknya. Kebaruan (*state of the art*) dari penelitian ini terletak pada penerapan kombinasi TF-IDF dan SVM secara spesifik pada domain ulasan kursus komersial lintas topik di platform UdeMy, yang diposisikan sebagai instrumen Business Intelligence untuk mendukung keputusan pengembangan kualitas kursus, suatu pendekatan yang belum banyak dieksplorasi pada riset-riset sentimen pendidikan sebelumnya yang umumnya berfokus pada institusi pendidikan formal atau aplikasi pembelajaran tunggal. Melalui temuan ini, instruktur maupun pengembang platform diharapkan memperoleh gambaran komprehensif mengenai ekspektasi pengguna, sekaligus mempermudah penyusunan skala prioritas dalam menyempurnakan fasilitas pengajaran.

2. Metodologi

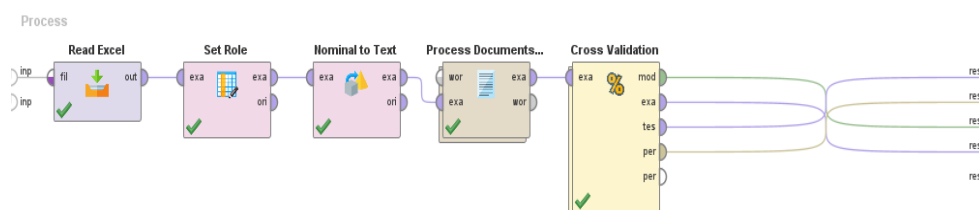
2.1. Prosedur Penelitian

Studi ini dieksekusi guna mengevaluasi kecenderungan opini dari pelanggan situs UdeMy dengan mengimplementasikan metode *Support Vector Machine* (SVM) yang dipadukan dengan pembobotan TF-IDF. Rangkaian penelitian secara substantif mencakup sepuluh tahapan yang saling berurutan, yaitu: (1) Pengumpulan Data, (2) Pembersihan Data, (3) Pelabelan Sentimen, (4) Preprocessing, (5) Ekstraksi Fitur, (6) Pembagian Data, (7) Pelatihan Model SVM, (8) Pengujian Model, (9) Evaluasi, dan (10) Analisis Hasil. Alur kerja tersebut diilustrasikan secara rinci pada Gambar 1.



Gambar 1. Diagram Alur Penelitian

Read Excel (memuat 50 data hasil pembersihan), Set Role (menetapkan atribut Sentimen sebagai label), Nominal to Text (mengonversi kolom ulasan menjadi tipe teks), Process Documents from Data (menjalankan tahap Preprocessing dan Ekstraksi Fitur sekaligus), dan Cross Validation (menjalankan tahap Pembagian Data, Pelatihan Model SVM, Pengujian Model, dan Evaluasi secara terintegrasi).



Gambar 2. Proses Utama pada RapidMiner

2.2. Pengumpulan Data

Populasi penelitian ini adalah seluruh ulasan pengguna pada platform UdeMy yang berkaitan dengan pengalaman belajar dan kualitas kursus. Pengumpulan data dilakukan melalui web scraping otomatis terhadap seluruh ulasan yang tersedia secara publik pada halaman kursus, tanpa penarikan sub-sampel pada tahap ini.

2.3. Pembersihan Data

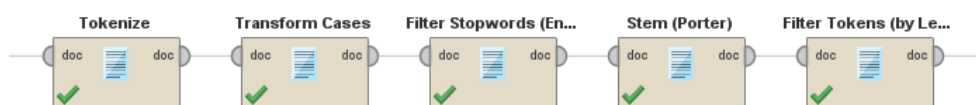
Terhadap ±500 data mentah hasil *scraping*, dilakukan pembuangan entri duplikat, ulasan kosong, dan ulasan berbahasa non-Inggris. Pemeriksaan distribusi *rating* pada data yang tersisa menunjukkan bahwa *rating* 1–2 (dasar label Negatif) jumlahnya jauh lebih sedikit dibandingkan *rating* 3–5 (dasar label Positif), sehingga diterapkan teknik *random undersampling* untuk menyeimbangkan kelas alih-alih menggunakan seluruh data secara langsung [13], [14]. Pendekatan ini dipilih karena distribusi kelas yang timpang berpotensi membuat model SVM bias terhadap kelas mayoritas dan menyulitkan interpretasi metrik evaluasi [13], [15].

2.4. Pelabelan Sentimen

Pelabelan sentimen pada penelitian ini tidak dilakukan secara manual maupun melalui pemrosesan otomatis oleh algoritma tertentu, melainkan murni mengadopsi sistem skor (*rating*) yang telah tersedia pada situs sumber data. Skala label yang digunakan bersifat biner, yaitu Positif dan Negatif. Ulasan dengan nilai bintang 3 sampai 5 ditetapkan sebagai sentimen Positif, sedangkan ulasan dengan nilai bintang 1 dan 2 ditetapkan sebagai sentimen Negatif. Pendekatan ini dipilih untuk menjaga objektivitas label karena bersumber langsung dari penilaian pengguna, sekaligus menghindari subjektivitas yang mungkin muncul apabila pelabelan dilakukan secara manual oleh peneliti.

2.5. Preprocessing

Pemrosesan teks ulasan dilakukan secara otomatis menggunakan modul *Process Documents from Data* pada perangkat lunak RapidMiner melalui lima tahapan berurutan. Tahap pertama, *Tokenization*, bekerja dengan memecah rangkaian kalimat ulasan menjadi token kata individual sekaligus menghapus karakter non-alfabet, angka, dan tanda baca secara otomatis. Selanjutnya, tahap *Transform Cases* mengonversi seluruh huruf dalam token menjadi bentuk huruf kecil (*lower case*) guna menyeragamkan variasi kapitalisasi teks. Pada tahap ketiga, operator *Filter Stopwords* (English) diterapkan untuk membuang kata-kata umum dan kata hubung dalam bahasa Inggris yang tidak memuat bobot sentimen diskriminatif. Tahap keempat dilanjutkan dengan *Stemming* menggunakan algoritma *Porter Stemmer* [16], yang berfungsi memotong imbuhan dan mengembalikan setiap token kata ke bentuk kata dasarnya. Rangkaian pemrosesan diakhiri dengan tahap *Filter Tokens (by Length)* yang menyaring token berdasarkan rentang panjang kata 3 hingga 25 karakter, sehingga token di luar rentang tersebut dihapus untuk mengeliminasi noise atau kata bentukan yang tidak valid. Rangkaian operator yang membentuk sub-proses ini pada RapidMiner ditunjukkan pada Gambar 3.



Gambar 3. Sub-proses Preprocessing (Process Documents from Data) pada RapidMiner

2.6. Ekstraksi Fitur

Dokumen ulasan direpresentasikan sebagai vektor numerik 398 dimensi menggunakan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF). Dengan metode ini, term khas yang jarang muncul di dokumen lain (seperti “*useless*”) mendapat bobot lebih tinggi dibandingkan term umum (seperti “*course*”). Perhitungan bobot menggunakan rumus berikut:

$$TF\text{-}IDF(t, d) = TF(t, d) \times \log(N/df(t)) \quad (1)$$

Keterangan: $N = 50$ (jumlah dokumen) dan $df(t)$ adalah jumlah dokumen yang memuat term t . Tahap ini menghasilkan matriks berukuran 50×398 sebagai masukan langsung untuk pelatihan model SVM.

2.7. Algoritma Support Vector Machine

Model klasifikasi dibangun menggunakan LibSVM dengan konfigurasi: tipe C-SVC, kernel Linear, parameter $C = 1.0$, dan epsilon = 0.001. Model menerima matriks TF-IDF 50×398 untuk mencari *hyperplane* yang memisahkan kelas Positif dan Negatif dengan *margin* maksimal dari titik data terdekat (*support vectors*), sesuai prinsip dasar SVM [17]. Kernel Linear dipilih karena representasi teks berdimensi tinggi umumnya dapat dipisahkan secara linear tanpa memerlukan komputasi kernel non-linear yang rumit. Secara matematis, keputusan klasifikasi model dirumuskan sebagai berikut:

$$f(x) = \text{sign}(w \cdot x + b) \quad (2)$$

dengan w dan b sebagai parameter hasil optimasi pada 45 data latih di setiap *fold*.

2.8. Pembagian Data

Penelitian ini menerapkan skema 10-Fold *Cross Validation* dengan membagi *dataset* 50 ulasan secara acak menjadi 10 *fold*. Pada setiap iterasi, 9 *fold* (45 data) digunakan sebagai data latih (*training*) dan 1 *fold* (5 data) sebagai data uji (*testing*). Proses ini diulang 10 kali sehingga semua data pernah diuji, guna memastikan pemanfaatan data yang optimal dan menghasilkan estimasi kinerja yang lebih stabil untuk *dataset* berukuran kecil. Pada RapidMiner (Gambar 4), operator *Cross Validation* memuat dua sub-proses: sisi *Training* (berisi operator SVM) dan sisi *Testing* (berisi operator *Apply Model* dan *Performance*).



Gambar 4. Sub-proses Training dan Testing dalam Cross Validation pada RapidMiner

2.9. Pengujian Model

Setiap model SVM yang telah dilatih pada Bagian 2.7 diterapkan pada data uji di *fold* yang bersesuaian untuk memperoleh label prediksi pada data yang belum pernah dipelajari model. Prosedur ini diulang untuk kesepuluh *fold* sehingga seluruh dokumen memperoleh tepat satu prediksi. Hasil prediksi dari seluruh *fold* kemudian diakumulasikan ke dalam satu confusion matrix gabungan yang menjadi dasar perhitungan metrik evaluasi pada Bagian 2.10.

2.10. Evaluasi Model

Kinerja model dihitung dari *confusion matrix* gabungan 10 *fold* (format ditunjukkan pada Tabel 6), yang memuat empat sel: *TP* (Positif diprediksi Positif), *TN* (Negatif diprediksi Negatif), *FP* (Negatif diprediksi Positif), dan *FN* (Positif diprediksi Negatif). Keempat sel tersebut dimasukkan ke rumus berikut untuk memperoleh *Accuracy*, *Precision*, *Recall*, dan *F1-Score* pada Tabel 7:

$$\text{Accuracy} = (TP + TN) / 50 \quad (3)$$

$$\text{Precision} = TP / (TP + FP) \quad (4)$$

$$\text{Recall} = TP / (TP + FN) \quad (5)$$

$$F1 - \text{Score} = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (6)$$

Perhitungan yang sama diterapkan pada kelas Negatif dengan menukar posisi TP/TN dan FP/FN, menghasilkan *Precision*, *Recall*, dan *F1-Score* khusus untuk kelas Negatif pada baris terpisah di Tabel 7.

3. Hasil dan Pembahasan

3.1. Hasil Pengumpulan Data

Proses web scraping terhadap halaman ulasan kursus Udemy dilakukan terhadap seluruh ulasan yang tersedia secara publik (populasi penelitian), dan menghasilkan luaran mentah sebanyak ± 500 ulasan berbahasa Inggris, yaitu keseluruhan ulasan yang berhasil terjaring, bukan sub-sampel. Data tersimpan dalam format CSV dengan dua kolom utama: teks ulasan (*content*) dan skor bintang (*rating* 1–5). Data mentah inilah yang menjadi masukan bagi tahap Pembersihan Data pada Bagian 3.2.

3.2. Hasil Pembersihan Data

Pemeriksaan terhadap ± 500 data mentah menunjukkan distribusi rating yang timpang: jumlah ulasan berating 1–2 (calon label Negatif) jauh lebih sedikit dibandingkan ulasan berating 3–5 (calon label Positif), dengan hanya 24 ulasan Negatif yang berhasil terjaring. Mengikuti teknik *random undersampling* pada Bagian 2.1 dan 2.2, seluruh 24 ulasan Negatif tersebut diikutsertakan seluruhnya, kemudian 26 ulasan Positif dipilih secara acak dari populasi Positif yang jauh lebih besar. Luaran konkret dari tahap ini adalah dataset bersih dan seimbang berjumlah 50 ulasan (26 Positif dan 24 Negatif), setelah entri duplikat, kosong, dan berbahasa non-Inggris dibuang.

3.3. Hasil Pelabelan Sentimen

Terhadap 50 data hasil pembersihan, label kelas ditetapkan berdasarkan skor bintang pada situs sumber: *rating* 3–5 diberi label Positif dan *rating* 1–2 diberi label Negatif. Rincian distribusi hasil pelabelan ini disajikan pada Tabel 1.

Tabel 1. Distribusi Label Sentimen

Sentimen	Jumlah
Positif	26
Negatif	24
Total	50

Berdasarkan Tabel 1, jumlah data Positif dan Negatif relatif seimbang (26 banding 24), sehingga risiko bias kelas pada proses klasifikasi SVM dapat diminimalkan.

3.4. Hasil Preprocessing

Kelima puluh teks ulasan diproses melalui modul *Process Documents from Data* pada RapidMiner dengan parameter operator sebagaimana ditunjukkan pada Tabel 2, dan menghasilkan 398 atribut (term) bersih dari 50 dokumen. Perlu ditekankan bahwa kolom *content* pada Tabel 3 bukan lagi kalimat gramatikal, melainkan kumpulan token/kata kunci (*bag-of-words*) yang tersisa setelah *tokenization*, *stopword removal*, dan *stemming*, yang ditampilkan berjejer sesuai urutan kemunculan aslinya sehingga sekilas menyerupai kalimat pendek. Contoh data setelah *pre-processing* ditunjukkan pada Tabel 3.

Tabel 2. Parameter *Preprocessing* pada RapidMiner

Operator	Parameter	Nilai
Tokenize	mode	non letters
Transform Cases	transform to	lower case
Filter Stopwords	-	English
Stem	-	Porter
Filter Tokens (by Length)	min/max chars	3 / 25

Tabel 3. Contoh Data Setelah *Pre-Processing*

No	Content	Sentimen
3	useless cours talk much basic	Negatif
36	love cours lot excersic need	Positif
...	...	...

3.5. Hasil Ekstraksi Fitur

Ke-398 term hasil preprocessing dikonversi menjadi matriks numerik berukuran 50×398 menggunakan pembobotan TF-IDF. Contoh hasil pembobotan ditunjukkan pada Tabel 4.

Tabel 4. Contoh Hasil Pembobotan TF-IDF

No	Sentimen	Term	Bobot TF-IDF
1	Negatif	walk	0,327
2	Negatif	bot	0,349
3	Negatif	useless	0,659
...	...	...	...

Kata-kata yang bersifat khusus dalam sebuah dokumen, misalnya “*useless*” maupun “*bot*”, mendapatkan nilai pembobotan yang lebih besar daripada kosakata lumrah yang sering ditemukan pada berbagai ulasan. Matriks 50×398 inilah yang menjadi masukan langsung bagi tahap Pembagian Data pada Bagian 3.6.

3.6. Hasil Pembagian Data

Matriks 50 dokumen hasil ekstraksi fitur dibagi ke dalam 10 subset (fold) berukuran hampir sama, masing-masing beranggotakan sekitar 5 dokumen. Pada setiap iterasi, 9 fold (45 dokumen) digunakan sebagai data latih dan 1 fold (5 dokumen) digunakan sebagai data uji, sehingga proses pelatihan dan pengujian pada Bagian 3.7 dan 3.8 diulang sebanyak 10 kali hingga seluruh 50 dokumen pernah menjadi data uji tepat satu kali.

3.7. Hasil Pelatihan Model SVM

Pada setiap fold, model SVM dilatih menggunakan 45 dokumen data latih dengan konfigurasi LibSVM sebagaimana ditunjukkan pada Tabel 5, sehingga diperoleh 10 model *hyperplane* pemisah kelas Positif dan Negatif (satu model per fold).

Tabel 5. Parameter SVM (LibSVM)

Parameter	Nilai
SVM type	C-SVC
Kernel type	Linear
C	1.0
Epsilon	0.001
Validasi	Cross Validation 10-fold

3.8. Hasil Pengujian Model

Setiap model hasil pelatihan pada Bagian 3.7 diuji terhadap 5 dokumen data uji pada fold yang bersesuaian. Hasil prediksi dari seluruh 10 fold pengujian kemudian diakumulasikan ke dalam satu *confusion matrix* gabungan yang mencakup keseluruhan 50 dokumen, sebagaimana ditunjukkan pada Tabel 6.

Tabel 6. *Confusion Matrix* Hasil Pengujian Model

	Prediksi Positif	Prediksi Negatif
Aktual Positif	22	4
Aktual Negatif	5	19

Dari 50 data yang diuji, model SVM mengklasifikasikan 22 dari 26 data Positif secara tepat (benar sebagai Positif) dan 19 dari 24 data Negatif secara tepat (benar sebagai Negatif). Sisanya, 4 data Positif justru diprediksi sebagai Negatif (*False Negative*), dan 5 data Negatif diprediksi sebagai Positif (*False Positive*). Dengan kata lain, dari 50 ulasan yang diuji, total 41 ulasan (22+19) diklasifikasikan secara benar dan 9 ulasan (4+5) diklasifikasikan keliru.

3.9. Hasil Evaluasi

Mengacu pada rumus di Bagian 2.10 dan nilai *confusion matrix* pada Tabel 6 (TP=22, TN=19, FP=5, FN=4), perhitungan konkretnya adalah sebagai berikut: $Accuracy = (22+19)/50 = 0,82$; untuk kelas Positif, $Precision = 22/(22+5) = 0,81$ dan $Recall = 22/(22+4) = 0,85$, sehingga $F1-Score = 2 \times (0,81 \times 0,85) / (0,81 + 0,85) = 0,83$; sedangkan untuk kelas Negatif, $Precision = 19/(19+4) = 0,83$ dan $Recall = 19/(19+5) = 0,79$, sehingga $F1-Score = 0,81$. Seluruh angka ini dirangkum pada Tabel 7.

Tabel 7. Hasil Evaluasi Kinerja SVM

	Precision	Recall	F1-Score	Support
Negatif	0,83	0,79	0,81	24
Positif	0,81	0,85	0,83	26
Accuracy			0,82	50
Macro avg	0,82	0,82	0,82	50
Weighted avg	0,82	0,82	0,82	50

Berdasarkan hasil pengujian tersebut, model mencapai tingkat akurasi 82% serta rata-rata f1-score 0,82. Pencapaian ini mengindikasikan bahwa metode *Support Vector Machine* (SVM) berbasis ekstraksi fitur TF-IDF memiliki keandalan yang memadai untuk memilah ulasan pengguna platform UdeMy menjadi kelompok sentimen positif maupun negatif.

3.10. Analisis Hasil Sentimen untuk Peningkatan Kursus

Guna mengamati distribusi topik di tiap kelompok sentimen, penelitian ini menghitung jumlah dokumen yang memuat masing-masing term (*document frequency*) pada matriks TF-IDF hasil ekstraksi fitur, dipilah berdasarkan label sentimen aktualnya. Term “cours” (*course*) dikecualikan dari daftar karena muncul hampir di seluruh ulasan pada kedua kelas sehingga tidak diskriminatif. Daftar kosakata yang paling mendominasi di setiap kategori klasifikasi disajikan secara rinci pada Tabel 8.

Tabel 8. Kata dengan Frekuensi Tertinggi pada Ulasan Positif dan Negatif

No	Kata Dominan Positif (frekuensi)	Kata Dominan Negatif (frekuensi)
1	learn (12)	time (8)
2	angela (10)	project (7)
3	great (8)	beginner (5)
4	practice (5)	instructor (5)
5	knowledge (5)	basic (5)
6	challenge (5)	waste (4)
7	enjoy (5)	lesson (4)
8	explain (5)	pycharm (4)
9	good (5)	thing (4)
10	lot (5)	work (4)

Merujuk pada Tabel 8, ulasan bernada positif mayoritas memuat kosakata seperti learn, angela (nama instruktur), great, serta practice, yang merepresentasikan tingginya kepuasan pengguna atas cara mengajar instruktur dan proses belajar praktikal pada kursus. Sebaliknya, area yang kerap menuai kritik pada ulasan negatif lebih sering ditandai oleh kemunculan istilah seperti time, project, beginner, serta instructor, yang menjadi sinyal bahwa pengguna masih menemui kendala menyangkut efisiensi alokasi waktu, beban proyek/tugas yang diberikan, kesesuaian materi bagi pemula, hingga gaya penyampaian instruktur yang dirasa belum optimal. Hasil pada Bagian 3.1 hingga 3.10 di atas menjadi dasar bagi pembahasan yang lebih mendalam pada Bagian 3.11.

3.11. Pembahasan

Merujuk pada ringkasan data evaluasi peserta UdeMy yang tersaji di Tabel 8, terlihat bahwa ulasan bernada positif mayoritas memuat kosakata semacam learn, angela (nama instruktur), great, serta practice. Pola ini konsisten merepresentasikan tingginya kepuasan pengguna atas cara mengajar instruktur dan proses belajar praktikal pada kursus. Sebaliknya, area yang kerap menuai kritik pada ulasan negatif lebih sering ditandai oleh kemunculan istilah macam time, project, beginner, serta instructor. Kosakata tersebut menjadi sinyal bahwa pengguna masih menemui kendala menyangkut efisiensi alokasi waktu, beban proyek/tugas yang diberikan, kesesuaian materi bagi pemula, hingga gaya penyampaian instruktur yang dirasa belum optimal.

Hasil evaluasi menunjukkan bahwa model SVM berbasis TF-IDF mencapai akurasi 82% dan rata-rata F1-score 0,82 (Tabel 7), yang mengonfirmasi bahwa model mampu memilah ulasan pengguna UdeMy ke dalam kelompok sentimen positif dan negatif dengan tingkat keandalan yang baik. Capaian ini secara langsung menjawab tujuan penelitian yang dirumuskan pada Bagian Pendahuluan, yaitu menganalisis opini konsumen terkait efektivitas kursus di UdeMy serta memetakan elemen edukasi yang memicu reaksi memuaskan maupun sebaliknya. Kemunculan kosakata dominan pada Tabel 8 memberikan gambaran konkret mengenai aspek-aspek kursus yang perlu dipertahankan maupun diperbaiki, sehingga tujuan penelitian untuk menghasilkan wawasan bagi instruktur dan pengembang platform tercapai.

Secara konseptual, capaian akurasi tersebut selaras dengan landasan teori yang menyatakan bahwa SVM efektif menangani data teks berdimensi tinggi melalui pembentukan *hyperplane* optimal antarkelas, serta bahwa representasi TF-IDF mampu menonjolkan istilah

yang secara khusus mencirikan suatu dokumen dibandingkan istilah umum [4], [9]. Nilai *precision* dan *recall* yang relatif seimbang antara kelas Positif (0,81 dan 0,85) dan Negatif (0,83 dan 0,79) pada Tabel 7 turut mengonfirmasi prinsip bahwa keseimbangan distribusi kelas pada tahap pengambilan sampel data (26 Positif dan 24 Negatif) berkontribusi menekan bias klasifikasi, sebagaimana menjadi dasar teoretis pemilihan skema sampling seimbang pada Bagian Metodologi.

Dibandingkan dengan penelitian Lubis dan Putri [9] yang menerapkan SVM pada sentimen mahasiswa terhadap *e-learning*, maupun studi Hanif dkk. [10] yang mengonfirmasi daya saing SVM dibandingkan *Decision Tree* dan *Random Forest* pada klasifikasi komentar mahasiswa, hasil penelitian ini menunjukkan tren yang konsisten, yaitu SVM tetap andal mengklasifikasikan sentimen berbasis teks pendek pada ranah pendidikan. Hasil ini juga sejalan dengan Thomas dan Rumaisa [18] yang membuktikan bahwa kombinasi TF-IDF dan SVM efektif mengklasifikasikan sentimen pada ulasan pengguna berbahasa alami, serta konsisten dengan Pranata dkk. [19] yang menegaskan bahwa pembobotan TF-IDF berkontribusi signifikan terhadap performa klasifikasi berbasis SVM pada domain berbeda (klasifikasi teks pemilu). Akan tetapi, tingkat akurasi pada penelitian ini (82%) relatif lebih rendah dibandingkan sejumlah studi lain yang menerapkan kombinasi serupa pada domain non-pendidikan, misalnya penelitian Aula dkk. [20] pada ulasan layanan ekspedisi maupun Que dkk. [21] pada ulasan transportasi daring yang melaporkan akurasi di atas 90%. Perbedaan ini kemungkinan besar disebabkan oleh jumlah sampel data yang jauh lebih kecil pada penelitian ini (50 ulasan) dibandingkan studi-studi tersebut yang umumnya menggunakan ribuan hingga puluhan ribu data, sehingga model memiliki keterbatasan dalam mempelajari variasi pola bahasa secara lebih luas. Perbandingan performa antar-algoritma turut memperkuat pemilihan SVM pada penelitian ini: Tarigan dkk. yang membandingkan SVM, *Naïve Bayes Classifier*, dan *Random Forest* pada analisis sentimen ulasan aplikasi Sirekap melaporkan bahwa SVM secara konsisten unggul dibandingkan *Naïve Bayes Classifier* [22]; temuan serupa dilaporkan Parameswari dkk. pada analisis sentimen pengguna *metaverse*, yang juga menunjukkan superioritas SVM atas *Naïve Bayes* [23]; sementara Nabarian dkk. membandingkan SVM dan *Naïve Bayes* pada klasifikasi sentimen evaluasi dosen dan mendapati bahwa SVM tetap lebih unggul dari *Naïve Bayes* meskipun keduanya masih kalah dibandingkan model berbasis IndoBERT [24]. Konsistensi keunggulan SVM atas *Naïve Bayes* lintas domain tersebut memperkuat rasionalisasi pemilihan SVM pada penelitian ini, sekaligus menjadi catatan bahwa pendekatan berbasis *pretrained language* model seperti IndoBERT berpotensi menjadi arah pengembangan lanjutan apabila volume data pelatihan berbahasa yang sesuai tersedia secara memadai. Temuan ini juga tidak sepenuhnya sejalan dengan studi Kurnia dkk. [8] dan Rahayu-Buditjahjanto [7] yang berfokus pada aspek administratif maupun tugas perkuliahan; penelitian ini justru memberikan temuan baru berupa pemetaan sentimen pada ranah kursus daring komersial lintas topik, yang secara spesifik diarahkan sebagai instrumen *Business Intelligence*, aspek yang belum banyak diulas pada riset-riset tersebut.

Kontribusi utama penelitian ini terletak pada penyediaan bukti empiris bahwa kombinasi TF-IDF dan SVM dapat diadaptasi ke ranah ulasan kursus daring komersial, memperluas cakupan penerapan analisis sentimen di bidang pendidikan yang selama ini banyak berfokus pada institusi pendidikan formal [4], [6], [7], [8]. Secara praktis, hasil analisis kosakata dominan (Tabel 8) memberikan kontribusi langsung bagi pengembangan manajemen mutu pembelajaran berbasis data (*data-driven quality management*), di mana penyedia kursus dapat memanfaatkan pola sentimen tersebut sebagai bahan evaluasi berbasis *Business Intelligence* [1], dan bukan sekadar hasil klasifikasi statistik semata.

Sembilan kesalahan klasifikasi pada Tabel 6 (4 *False Negative* dan 5 *False Positive*) juga memberi petunjuk mengenai batas kemampuan model pada dataset berukuran kecil ini. Empat ulasan Positif yang salah diprediksi sebagai Negatif kemungkinan besar memuat campuran kosakata yang tumpang tindih antar kelas, misalnya kata *python* dan *like* yang sama-sama muncul pada daftar kata dominan Positif maupun Negatif di Tabel 8; kondisi ini membuat bobot TF-IDF dari kata-kata tersebut kurang diskriminatif sehingga hyperplane SVM lebih sulit memisahkan kedua kelas pada dokumen yang bersangkutan. Pola serupa berlaku pada lima ulasan Negatif yang salah diprediksi sebagai Positif, yang diduga memuat kritik dengan nada yang relatif halus (misalnya keluhan disampaikan bersamaan dengan pujian pada bagian lain ulasan), sehingga secara leksikal lebih menyerupai pola ulasan Positif. Temuan ini menunjukkan bahwa keterbatasan kinerja model bukan semata-mata pada pilihan algoritma SVM, melainkan

juga pada kualitas pemisahan fitur pada dataset yang berukuran kecil dan mengandung kosakata tumpang tindih antar kelas.

Meskipun demikian, penelitian ini memiliki sejumlah keterbatasan. Pertama, ukuran sampel yang digunakan tergolong kecil (50 ulasan) akibat keterbatasan distribusi kelas pada data hasil *scraping*, sehingga generalisasi hasil pada populasi ulasan Udemy secara keseluruhan perlu dilakukan secara hati-hati. Sebagaimana ditegaskan oleh Agustian dkk. [25], performa model klasifikasi sentimen berbasis teks sangat rentan terhadap keterbatasan ukuran data, dan penambahan volume maupun teknik augmentasi data dapat menjadi strategi untuk meningkatkan generalisasi model pada penelitian selanjutnya. Kedua, penelitian ini hanya menggunakan skema pelabelan biner (positif dan negatif) berdasarkan kategori rating situs sumber, tanpa mempertimbangkan kelas netral maupun intensitas sentimen secara lebih granular. Ketiga, penelitian ini hanya menerapkan kernel Linear pada algoritma SVM tanpa membandingkan dengan varian kernel lain seperti RBF atau polynomial, maupun algoritma klasifikasi pembandingan seperti *Naïve Bayes* atau *Random Forest*, sehingga belum dapat dipastikan apakah performa yang diperoleh merupakan hasil yang paling optimal. Oleh karena itu, penelitian lanjutan disarankan untuk memperluas volume dan keragaman dataset, mengeksplorasi skema pelabelan multikelas, serta melakukan uji banding kernel dan algoritma klasifikasi lain guna memperoleh kesimpulan yang lebih komprehensif mengenai efektivitas metode dalam mengevaluasi pembelajaran pada platform kursus daring.

Berpijak pada wawasan tersebut, langkah strategis yang dapat dieksekusi guna mendongkrak mutu kursus mencakup dua arah utama. Penyedia layanan perlu menjaga standar pengajaran dan metode belajar praktikal yang sudah terbukti memuaskan, sekaligus segera mengevaluasi kedalaman substansi materi, penyesuaian durasi kelas, dan efektivitas komunikasi pengajar agar pengembangannya lebih selaras dengan ekspektasi peserta.

4. Simpulan

Studi ini difokuskan untuk mengevaluasi opini pengguna mengenai evaluasi pembelajaran di platform Udemy. Metode yang diterapkan melibatkan algoritma *Support Vector Machine* (SVM) yang dipadukan bersama teknik ekstraksi fitur *Term Frequency-Inverse Document Frequency* (TF-IDF). Melalui tahapan pengujian dengan skema *10-Fold Cross Validation*, model klasifikasi ini mampu mencetak tingkat akurasi hingga 82% serta *F1-score* rata-rata di angka 0,82. Angka tersebut membuktikan bahwa perpaduan TF-IDF dan SVM tergolong andal untuk memetakan sentimen ulasan peserta pada *dataset* yang digunakan.

Temuan riset mengungkapkan bahwa ulasan bernada positif mayoritas memuat kosakata seperti *learn*, *angela* (nama instruktur), *great*, serta *practice*, yang merepresentasikan tingkat kepuasan peserta atas cara mengajar instruktur dan proses belajar praktikal pada kursus. Sebaliknya, keluhan pada ulasan negatif lebih sering ditandai oleh kemunculan istilah macam *time*, *project*, *beginner*, serta *instructor*. Kosakata tersebut mengindikasikan adanya isu seputar manajemen waktu, beban proyek, kesesuaian materi bagi pemula, hingga cara pengajar menyampaikan materi. Hal ini menegaskan bahwa penerapan analisis sentimen tidak sebatas memetakan tren opini, melainkan turut berperan penting dalam mendeteksi area spesifik yang butuh dievaluasi atau dipertahankan kualitasnya.

Mengacu pada temuan ini, pihak penyedia layanan kursus dapat menjadikan hasil analisis sentimen sebagai rujukan untuk menjaga standar aspek pembelajaran yang sudah dinilai memuaskan, sekaligus membenahi poin-poin yang masih menuai kritik. Ke depannya, penelitian lanjutan sangat direkomendasikan untuk memperluas cakupan volume *dataset* serta melakukan komparasi performa antara SVM dengan metode klasifikasi lainnya demi mencapai kesimpulan yang lebih utuh.

Daftar Referensi

- [1] M. Wijana, R. Cahya Gumelar, R. D. Supriatman, and Y. Muhyidin, "Aplikasi Sistem Pendukung Keputusan Menentukan Siswa Terbaik Menggunakan Metode SAW," *INTERNAL (Information System Journal)*, vol. 7, no. 1, pp. 18–29, Jun. 2024, doi: 10.32627/internal.v7i1.973.
- [2] D. Gitawijaya, A. Siswandi, and I. Afriantoro, "Analisis Sentimen Indihome Dan Starlink Menggunakan Metode *Naïve Bayes* Bagi Pt. Sarimelati Kencana, Tbk Dalam Pengambilan Keputusan," *DINAMIK*, vol. 30, no. 2, pp. 339–350, 2025, doi: 10.35315/dinamik.v30i2.10252.

- [3] R. Cahyana, A. Nurul Hakim, F. Al-Husein, and D. Wildan Munawar, "Tinjauan Persepsi dan Analisis Sentimen Mahasiswa Dalam Implementasi Program Merdeka Belajar Kampus Merdeka," *Journal of Digital Literacy and Volunteering*, vol. 4, no. 1, pp. 15–19, Jan. 2026, doi: 10.57119/litdig.v4i1.193.
- [4] A. I. Kamil, O. N. Pratiwi, and D. Witasryah, "Analisis Sentimen dan Pemodelan Topik terhadap Aplikasi Pembelajaran Online pada Platform Google Play," *JUPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, vol. 10, no. 2, pp. 836–849, Mar. 2025, doi: 10.29100/jupi.v10i2.6023.
- [5] A. Saninah, W. Prihartono, and C. L. Rohmat, "Analisis Sentimen Ulasan Pengguna Terhadap Aplikasi Pembelajaran Berbahasa Duolingo Menggunakan Algoritma Naïve Baiyes Classifier," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 1, pp. 619–628, Jan. 2025, doi: 10.23960/jitet.v13i1.5691.
- [6] M. Hilmy Maulana, R. A. Sanchia, and D. A. Prasetya, "Persepsi Publik terhadap Pendidikan Tinggi dalam Komentar YouTube Mata Najwa: Analisis Sentimen dengan Algoritma Support Vector Machine," *Seminar Nasional Amikom Surakarta (SEMNASA)*, Nov. 2025.
- [7] E. T. Rahayu and I. G. P. A. Buditjahjanto, "Analisis Sentimen Mahasiswa terhadap Tugas-Tugas Kuliah dengan Menggunakan Metode K-Nearest Neighbors," *SENTRI: Jurnal Riset Ilmiah*, vol. 5, no. 1, pp. 465–479, Jan. 2026, doi: 10.55681/sentri.v5i1.5425.
- [8] Nicky Dwi Kurnia, B. Kholifah, F. W. Nur Ani, and A. F. N. Nafii', "Analisis Sentimen Multi-Kelas untuk Menilai Kepuasan Mahasiswa terhadap Aplikasi Manajemen Akademik Berbasis Web di Lingkungan Perguruan Tinggi," *Jurnal Komputer Teknologi Informasi Sistem Informasi (JUKTISI)*, vol. 4, no. 1, pp. 194–197, Jun. 2025, doi: 10.62712/juktisi.v4i1.363.
- [9] A. S. Lubis and R. A. Putri, "Analisis Sentimen Mahasiswa Terhadap Penggunaan E-Learning dengan Algoritma Support Vector Machine (SVM)," *Decode: Jurnal Pendidikan Teknologi Informasi*, vol. 5, no. 2, pp. 778–788, Jul. 2025, doi: 10.51454/decode.v5i2.1247.
- [10] K. H. Hanif, N. R. Muntari, D. Harto, and S. Wiranata, "Perbandingan Analisis Sentimen Komentar Mahasiswa Prodi Teknik Komputer Menggunakan Algoritma Decision Tree, Support Vector Machine (SVM), dan Random Forest," *INSECT*, vol. 12, no. 01, pp. 24–33, 2026, doi: 10.33506/insect.v12i01.5144.
- [11] Mutiara Sintia Dewi and A. H. Hasugian, "Penerapan Support Vector Machine Untuk Analisis Sentimen Pada Tanggapan Masyarakat Di Media Sosial Terhadap Program Makan Siang Gratis," *Rabit : Jurnal Teknologi dan Sistem Informasi Univrab*, vol. 10, no. 2, pp. 911–922, Jul. 2025, doi: 10.36341/rabit.v10i2.6425.
- [12] D. Irawan, H. Budianto, and D. Puteria Wati, "Analisis Sentimen Mahasiswa Program Studi Sistem Informasi UNIKU terhadap Pelayanan Akademik dan Administrasi Menggunakan Metode Naive Bayes Classifier," *Buffer Informatika*, vol. 12, no. 1, pp. 74–79, Apr. 2025, doi: 10.25134/buffer.v12i1.573.
- [13] N. Chamidah, D. Widiyanto, H. B. Seta, and A. A. Aziz, "The Impact of Oversampling and Undersampling on Aspect-Based Sentiment Analysis of Indramayu Tourism Using Logistic Regression," *Revue d'Intelligence Artificielle*, vol. 38, no. 3, pp. 795–804, Jun. 2024, doi: 10.18280/ria.380306.
- [14] S. F. Taskiran, B. Turkoglu, E. Kaya, and T. Asuroglu, "A comprehensive evaluation of oversampling techniques for enhancing text classification performance," *Sci. Rep.*, vol. 15, no. 1, pp. 1–20, Dec. 2025, doi: 10.1038/s41598-025-05791-7.
- [15] A. Salehi and M. Khedmati, "Hybrid clustering strategies for effective oversampling and undersampling in multiclass classification," *Sci. Rep.*, vol. 15, no. 1, pp. 1–20, Dec. 2025, doi: 10.1038/s41598-024-84786-2.
- [16] A. S. Rizki, A. Tjahyanto, and R. Trialih, "Comparison of stemming algorithms on Indonesian text processing," *Telkomnika (Telecommunication Computing Electronics and Control)*, vol. 17, no. 1, pp. 95–102, Feb. 2019, doi: 10.12928/Telkomnika.v17i1.10183.
- [17] G. Khyathi, K. P. Indumathi, H. A. Jumana, F. J. M. Lisa, S. Sibyl, and G. Krishnaprakash, "Support Vector Machines: A Literature Review on Their Application in Analyzing Mass Data for Public Health," *Cureus*, vol. 17, no. 1, pp. 1–6, Jan. 2025, doi: 10.7759/cureus.77169.

- [18] V. W. D. Thomas and F. Rumaisa, "Analisis Sentimen Ulasan Hotel Bahasa Indonesia Menggunakan Support Vector Machine dan TF-IDF," *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 6, no. 3, p. 1767, Jul. 2022, doi: 10.30865/mib.v6i3.4218.
- [19] R. A. Pranata, Rudiman, and N. A. Verdikha, "Metode Pembobotan Tf-Idf Untuk Klasifikasi Teks Quick Count Pemilihan Wakil Presiden Indonesia 2024 Pada X Twitter Dengan Metode SVM," *Jurnal Keilmuan dan Aplikasi Bidang Teknik Informatika*, vol. 18, no. 2, p. 126, Jul. 2024, doi: 10.47111/JTI.
- [20] N. Aula, M. Ula, and L. Rosnita, "Analisis Sentimen Review Customer Terhadap Perusahaan Ekspedisi Jne, J&T Express Dan Pos Indonesia Menggunakan Metode Support Vector Machine (SVM)," *Journal of Informatics and Computer Science*, vol. 9, no. 1, pp. 81–86, Apr. 2023, doi: 10.33143/jics.v9i1.2947.
- [21] V. K. S. Que, A. Iriani, and H. D. Purnomo, "Analisis Sentimen Transportasi Online Menggunakan Support Vector Machine Berbasis Particle Swarm Optimization (Online Transportation Sentiment Analysis Using Support Vector Machine Based on Particle Swarm Optimization)," *Jurnal Nasional Teknik Elektro dan Teknologi Informasi* |, vol. 9, no. 2, May 2020, Accessed: Aug. 04, 2026. [Online]. Available: https://d1wqtxts1xzle7.cloudfront.net/89252955/116-libre.pdf?1659589481=&response-content-disposition=inline%3B+filename%3DAnalisis_Sentimen_Transportasi_Online_Me.pdf&Expires=1785861706&Signature=ZSY7kv69y2yxjvTfNQQAkrtrfR9v18Gf9cUFccniRcc5~BKT6phbP~NxpOBfOMTyw0Z-UKjPJy3fdGp9lmimBSgW~SGUYzPRe~10IHAvUGWmQmlmwhMxcrYkeU~KSq~6mc~qLhwd02d605m6B7VUuXb5n-17tc7MdYLnTxpfMz9vubr08BTr~~uKQMMvRwDua3eVobnPbNbB-X76xUviLZIfYQKIPa-hrEL7A7ztiUqdNSKUktUbiYzxUGTR9VywYtZHFuNG3pfHqG5VhXrRPaFCxPjyLryOVQ4BP2yTVeY9p4C1ToqXf3KuTQiL4sGmrNu4SeUW~MXfvIL32X8Hw__&Key-Pair-Id=APKAJLOHF5GGSLRBV4ZA
- [22] D. A. Tarigan, Z. Situmorang, and R. Rosnelly, "Analisis Sentimen Aplikasi Playstore Sirekap 2024 Pasca Pilpres Dengan Perbandingan Metode Support Vector Machine (Svm), Naïve Bayes Classifier Dan Random Forest," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIIK)*, vol. 11, no. 3, pp. 661–670, Jun. 2025, doi: 10.25126/jtiik.2025129608.
- [23] S. D. Parameswari *et al.*, "Studi Perbandingan Naïve Bayes dan Support Vector Machine (SVM) dalam Analisis Sentimen Pengguna Metaverse," *Jurnal Teknologi dan Manajemen Industri Terapan (JTMIT)*, vol. 4, no. 3, pp. 1059–1065, Sep. 2025, doi: 10.55826/jtmit.v4i3.1122.
- [24] T. Nabarian, M. H. Syamila, S. El Farisi, and A. D. Saputro, "Efektivitas IndoBERT pada Klasifikasi Sentimen Evaluasi Dosen: Studi Komparatif Support Vector Machine dan Naive Bayes," *Jurnal Teknologi Informasi dan Multimedia*, vol. 8, no. 2, pp. 281–292, Apr. 2026, doi: 10.35746/jtim.v8i2.984.
- [25] S. Agustian, M. Irfan Syah, N. Fatiara, and R. Abdillah, "New Directions in Text Classification Research: Maximizing The Performance of Sentiment Classification from Limited Data," *ArXiv*, Jul. 2024, doi: 10.48550/arXiv.2407.05627.