

Implementasi Metode *Latent Dirichlet Allocation* Pada Analisis Sentimen Berbasis Aspek Studi Kasus Aplikasi Spotify *(Implementation of the Latent Dirichlet Allocation Method in Aspect-Based Sentiment Analysis: A Case Study on the Spotify Application)*

Ester Gabriella Runtu^{1*}, Sunneng Sandino Berutu², Jatmika³

^{1,2,3} Program Studi Informatika, Fakultas Sains dan Komputer, Universitas Kristen Immanuel Yogyakarta
Jln. Ukrim km 11.1, Purwomartani, Kalasan, Sleman, D.I Yogyakarta, 55571, Indonesia

Corresponding author: estergabriella18@student.ukrimuniversity.ac.id

(Submitted: May 2025 / Reviewed: June 2025 / Accepted: June 2025)

ABSTRACT

Aspect-Based Sentiment Analysis (ABSA) is an important approach in understanding user opinions regarding specific features or services of a digital product. This study implements the Latent Dirichlet Allocation (LDA) method to identify the main aspects found in user reviews of the Spotify application, as well as to classify user sentiment toward each aspect. The data were obtained from 10,000 Indonesian-language reviews on the Google Play Store using a web scraping technique. The analysis process involved stages of text preprocessing, topic extraction using LDA, sentiment labeling using the InSet Lexicon, and classification evaluation using the Random Forest algorithm. The results show that there are three main aspects frequently discussed by users, namely Accessibility, Features, and User Experience, with a sentiment distribution that shows a dominance of negative sentiment. The implementation of this method demonstrates high accuracy in both aspect and sentiment classification. These findings contribute to providing data-driven insights that can be used by Spotify developers to improve the quality of user experience in a more targeted manner.

Keywords:

Aspect-Based Sentiment Analysis (ABSA); Implementation of the Latent Dirichlet Allocation (LDA) Method; Spotify Application;

ABSTRAK

Analisis Sentimen Berbasis Aspek (*Aspect-Based Sentiment Analysis/ABSA*) merupakan pendekatan penting dalam memahami opini pengguna terhadap fitur atau layanan tertentu dari suatu produk digital. Penelitian ini mengimplementasikan metode *Latent Dirichlet Allocation* (LDA) untuk mengidentifikasi aspek – aspek utama dalam ulasan pengguna aplikasi *Spotify*, serta mengklasifikasi sentimen pengguna terhadap setiap aspek. Data diperoleh dari 10.000 ulasan berbahasa Indonesia di *Google Play Store* menggunakan Teknik *web scraping*. Proses analisis melibatkan tahap *text preprocessing*, ekstraksi topik dengan LDA, pelabelan sentimen menggunakan kamus *InSet Lexicon*, serta evaluasi klasifikasi menggunakan algoritma *Random Forest*. Hasil penelitian menunjukkan bahwa terdapat tiga aspek utama yang sering dibahas oleh pengguna yaitu aspek *Accessibility*, *Features*, dan *User Experience*, dengan persebaran sentiment yang menunjukkan dominasi sentiment negative. Implementasi metode ini menunjukkan akurasi tinggi dalam klasifikasi aspek dan sentiment. Temuan ini memberikan kontribusi dalam menyediakan wawasan berbasis data yang dapat digunakan oleh pengembang *Spotify* untuk meningkatkan kualitas pengalaman pengguna secara lebih terarah.

Kata kunci:

Analisis Sentimen Berbasis Aspek (ABSA); Aplikasi Spotify; Implementasi Metode *Latent Dirichlet Allocation* (LDA);

1. Pendahuluan

Di era digital yang berkembang sangat pesat ini tidak dipungkiri kebutuhan teknologi yang semakin meningkat, terutama dalam dunia hiburan. Salah satu dunia hiburan yang berkembang yaitu hiburan music, ada banyak platform musik yang memungkinkan kita

untuk dapat mendengarkan music secara daring dan dapat kita streaming. Ada beberapa platform musik streaming yang ada salah satu contoh platform musik streaming yaitu Spotify, sebuah layanan musik digital, podcast, dan video yang memberimu akses ke jutaan lagu dan konten lain dari kreator di seluruh dunia. Per awal 2025, aplikasi Spotify telah diunduh lebih dari satu miliar kali dengan sekitar 30 juta ulasan pengguna di *Google Play Store*.

Tingginya interaksi pengguna pada aplikasi ini melalui ulasan digital memberikan potensi besar untuk memahami pengalaman dan persepsi pengguna terhadap aplikasi Spotify secara mendalam. Salah satu pendekatan yang relevan untuk tujuan tersebut adalah dengan Analisis Sentimen Berbasis Aspek (*Aspect-Based Sentiment Analysis/ABSA*) yaitu analisis sentiment yang memfokuskan pada aspek tertentu dalam sebuah teks. ABSA mampu mengidentifikasi sentien terhadap aspek tertentu dalam sebuah layanan.

Pada penelitian ini menggunakan metode *Latent Dirichlet Allocation* (LDA) sebagai pendekatan topic modelling untuk mengidentifikasi aspek – aspek utama dari ulasan pengguna. LDA dipilih karena kemampuannya dalam menangkap distribusi topik secara tidak terawasi serta kemudahan interpretasi terhadap Kumpulan kata – kata yang membentuk satu topik atau aspek. Walaupun terdapat metode alternatif lain seperti NMF (*Non-negative Matrix Factorization*) dan *BERTopic* (yang berbasis transformer), LDA masih relevan karena efisien dan dapat diintegrasikan dengan *pipeline text preprocessing* yang mapan, serta sudah terbukti efektif dalam studi – studi sebelumnya yang sejenisnya. Namun hasil dari analisis LDA tidak cukup jika tidak dilakukan bersama dengan pengukuran sentiment pengguna terhadap tiap aspek. Oleh karena itu tahap berikutnya akan dilakukan pelabelan sentiment dengan menggunakan kamus *InSet Lexicon* yang merupakan leksikon Bahasa Indonesia yang memuat ribuan kata dengan polaritas positif dan negatif. Dan pada tahap evaluasi performa klasifikasi aspek dan sentiment digunakan algoritma *Random Forest* yang dikombinasikan dengan TF-IDF dan Teknik oversampling dengan SMOTE untuk mengatasi data yang tidak seimbang.

Berdasarkan latar belakang diatas, penelitian ini bertujuan untuk mengidentifikasi aspek – aspek utama yang muncul dari ulasan pengguna Spotify menggunakan metode LDA, mengklasifikasikan sentimen pengguna terhadap masing – masing aspek menggunakan leksikon Bahasa Indonesia, dan mengevaluasi performa klasifikasi menggunakan algoritma *Random Forest* dan *K-Fold Cross Validation*. Diharapkan hasil dari penelitian ini dapat memberikan wawasan berbasis data untuk pengembangan aplikasi Spotify dalam meningkatkan kualitas layanan dan pengalaman pengguna secara lebih terfokus.

2. Metodologi

Pada tahap ini dilakukan pengumpulan informasi dan perancangan - perancangan sistem yang kemudian dilakukan analisis sesuai dengan data yang sudah didapat.

2.1. Material

2.1.1 Spotify

Menurut (Spotify.com, 2018) “Spotify adalah aplikasi music streaming asal Swedia yang memberikan akses untuk mendengarkan jutaan lagu dan podcast secara streaming ke pengguna”. (Karyono et al., 2019)

2.1.2 ABSA (*Aspect Based Sentiment Analysis*)

Analisis Sentimen Berbasis Aspek, atau ABSA adalah suatu jenis analisis sentimen yang bertujuan untuk mengetahui sentimen dalam setiap aspek yang ditentukan. ABSA memproses informasi pada tingkat subkalimat atau tingkat aspek. Dalam beberapa penelitian, proses dalam ABSA terbagi menjadi dua tugas yaitu tugas ekstraksi aspek dan memperkirakan polaritas/rating. (Aritonang et al., 2022)

2.1.3 Web Scraping

Web Scraping adalah metode yang digunakan untuk mengekstrak informasi dari website. *Web scraping* telah menjadi solusi yang populer untuk mengekstrak metadata dari halaman web. Banyaknya data yang dipublikasikan melalui website membuat menarik para peneliti tertarik untuk mengumpulkan dan mengolah menjadi informasi. (Purnomo, 2022)

2.1.4 Text Mining

Text Mining ialah sebagai teknik mengekstrak informasi dari sekumpulan data tidak terstruktur berkualitas tinggi serta diperoleh data – data masalah – masalah dalam teks topik tertentu. *Text Mining* dapat menemukan informasi penting dari sumber data dengan mengidentifikasi dan memeriksa pola tertentu. (Mustofa & Novita, 2022)

2.1.5 API Google – Play – Scraper

API Google-PlayScraper 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. (Larasati et al., 2022)

2.1.6 Text Cleaning

Text Cleaning akan menghilangkan tanda baca, karakter, huruf tunggal yang dianggap tidak memiliki arti dalam pemrosesan kata sehingga menghasilkan kalimat yang bersih. (Fauzi et al., 2021)

2.1.7 Text Preprocessing

Preprocessing menjadi tahap awal dalam klasifikasi teks untuk mempersiapkan data teks sebelum digunakan pada proses lainnya. Pada tahap ini akan mengubah data teks menjadi bentuk yang lebih baik sehingga menghasilkan informasi teks dengan kualitas yang baik dan siap digunakan pada proses selanjutnya. (Khairunnisa et al., 2021)

2.1.8 Latent Dirichlet Allocation

Pengertian *Latent Dirichlet Allocation* (LDA) merupakan sebuah metode statistika yang digunakan sebagai model untuk menganalisis suatu dokumen. LDA berusaha untuk melihat dokumen dengan cara mundur untuk menemukan satu set topik yang mungkin telah dikoleksi. LDA merepresentasikan dokumen dengan berbagai topik yang dibuat berdasarkan probabilitas tertentu. Probabilitas topik, merepresentasikan kejelasan dari suatu dokumen. (Hikmah et al., 2020)

2.1.9 Opinion Extraction

Menurut penelitian (Rozi, et al., 2012) “Pada proses ini setiap dokumen akan dipecah menjadi beberapa opini untuk mendapatkan aspek dan sentimen dari setiap ulasan. Untuk mendapatkan opini dari suatu dokumen dilakukan proses *POS tagging* yang akan memberikan kelas kata (tag) terlebih dahulu pada setiap kata dalam dokumen. Dari hasil kelas kata yang telah didapatkan tersebut, akan dilakukan chunking yaitu memecah sebuah informasi menjadi potongan – potongan kecil untuk mendapatkan hasil opini. Tag yang akan diambil dan diidentifikasi sebagai opini dalam suatu ulasan adalah *tak noun-phrase* (NP) yang mendefinisikan objek sebagai opini dalam suatu ulasan”. (Rahma Yustihan & Pandu Adikara, 2021)

2.1.10 Weighting Word

Weighting Word atau pembobotan kata merupakan mekanisme untuk memberikan skor pada frekuensi kemunculan suatu kata dalam dokumen teks. Salah satu metode yang populer untuk pembobotan kata adalah TF-IDF (*Term frequency-Inverse Document Frequency*). *Term frequency-Inverse Document Frequency* merupakan metode pembobotan yang menggabungkan dua konsep yaitu *Term Frequency* dan *Document Frequency*. (Yulita, 2021)

2.1.11 InSet Lexicon

InSet Lexicon merupakan kamus *lexicon* yang berisi 3.609 kata positif dan 6.609 kata negatif dalam Bahasa Indonesia dengan nilai *polarity* berkisar -5 sampai +5. (Fauziah, 2023)

2.1.12 Splitting Data

Tahapan ini dilakukan dengan cara memisahkan data menjadi data latih (*data training*) dan data uji (*data testing*). (Darusalam et al., 2024)

2.1.13 Random Forest

Menurut penelitian (Rohman, Purwanto, & Santoso, 2018) “*Random Forest* adalah metode *machine learning* yang terdiri dari gabungan *Decision Tree* untuk dilakukan klasifikasi dimana untuk memperoleh keputusan akhir dilakukan voting majority”. (Larasati et al., 2022)

2.1.14 SMOTE (Synthetic Minority Oversampling Technique)

SMOTE (*Synthetic Minority Oversampling Technique*) adalah teknik yang digunakan untuk mengatasi masalah ketidakseimbangan data dalam pembelajaran mesin. (Sidiq et al., 2025)

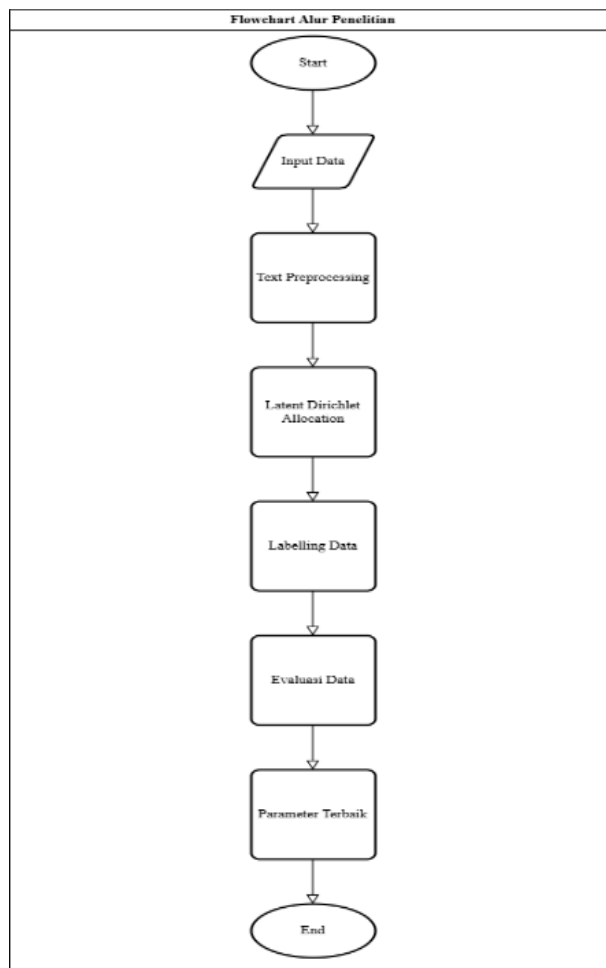
2.1.15 K-Fold Cross Validation

Menurut penelitian (Tempola, Muhammad, & Khairan, 2018) “*K-Fold Cross Validation* merupakan model validasi silang yang digunakan untuk memprediksi model dan memprediksi keakuratan model. Tujuan dilakukannya *K-Fold Cross Validation* yaitu untuk menghilangkan data yang bias”. (Larasati et al., 2022)

2.2. Prosedur

2.2.1 Flowchart Sistem

Gambaran umum sistem yang akan dibangun dapat dilihat pada Gambar 1 berikut.



Gambar 1. *Flowchart Alur Penelitian*

2.2.2 Perancangan Data

Data akan disimpan di *Google Drive* dan dikumpulkan melalui Teknik scraping menggunakan *Google Play Scraper*. Sebanyak 10.000 ulasan pengguna Spotify dalam Bahasa Indonesia dikumpulkan dengan pengaturan *lang='id'*. Data yang diambil mencakup *username*, *rating*, tanggal, dan ulasan, lalu disortir menjadi content dan rating sebelum disimpan sebagai file CSV di *Google Drive*.

2.2.3 Text Preprocessing

Dalam tahap *Text Preprocessing* akan melalui beberapa tahap diantaranya tahap *text cleaning*, *lowercase*, tokenisasi, *stopword*, dan *stemming*. Kemudian pada tahap ini data akan diseleksi agar menjadi lebih terstruktur, sehingga data yang dihasilkan akan lebih efektif dan akurat.

2.2.4 LDA (Latent Dirichlet Allocation)

Pada proses LDA (*Latent Dirichlet Allocation*) ada beberapa tahap yang akan dilakukan yaitu mempersiapkan model dan kamus LDA, proses koherensi untuk menentukan nilai num topik dan dominan topiknya, proses pyLDAvis untuk memvisualisasikan jumlah topik yang dominan, dan yang terakhir menentukan aspek untuk setiap topiknya.

2.2.5 Labelling Data

Tahap labelling data ini akan menggunakan *InSet Lexicon* sebagai alat untuk mencari sentimennya, karena data ulasan menggunakan bahasa Indonesia dan *InSet Lexicon* digunakan

untuk data berbahasa Indonesia. Kemudian akan dihitung jumlah setiap sentimennya dan disimpan dalam file csv.

2.2.6 Visualisasi Data

Proses Visualisasi data menggunakan diagram batang dan *wordcloud* untuk menampilkan atau memvisualisasikan hasil dari *Latent Dirichlet Allocation* (LDA) dan *labelling*.

2.2.7 Evaluasi Data

Random Forest digunakan untuk klasifikasi sentiment dan aspek. Proses ini menerapkan *TF-IDF*, *Stratified Split* (80% train, 20% test), dan *SMOTE* untuk menangani data tidak seimbang. Hasil evaluasi divisualisasikan dengan confusion matrix dan diagram batang.

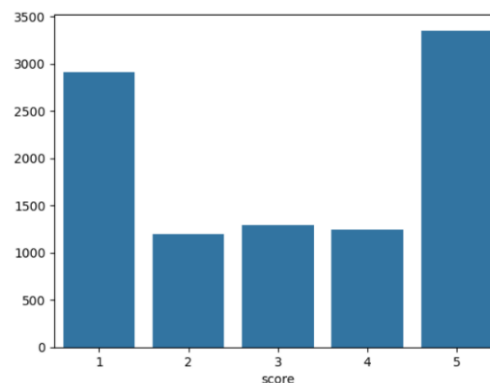
2.2.8 Parameter Terbaik

Untuk mencari parameter atau model terbaik akan menggunakan *K-Fold Cross Validation*. Dalam proses ini akan menghitung nilai skor setiap Fold nya yang terdiri dari 5 Fold dengan menggunakan *Stratified K-Fold Cross Validation* untuk data yang tidak seimbang. Hasil dari proses ini akan menjadi parameter atau model terbaik untuk penelitian ini.

3. Hasil dan Pembahasan

3.1. Hasil Pengumpulan Data

Pengujian ini menggunakan 10.000 ulasan *Spotify* dari *Google Play Store* yang di-*scraping* dengan lang='id' untuk bahasa Indonesia. Data mencakup *content* dan *score* (rating), dengan distribusi: *score* 1 (2.914 ulasan), *score* 2 (1.200), *score* 3 (1.290), *score* 4 (1.243), dan *score* 5 (3.353). Hasil dapat dilihat pada Gambar 2 berikut.



Gambar 2. Visualisasi Jumlah Setiap Score

3.2. Hasil Text Processing

Text preprocessing bertujuan membersihkan data mentah agar siap untuk analisis yang akurat. Proses ini menggunakan spaCy dan NLTK untuk pemrosesan Bahasa alami serta kamus Sastrawi untuk Bahasa Indonesia. Terdiri dari 5 tahap yaitu text cleaning, lower case, tokenisasi (spaCy), stopword (NLTK dan Sastrawi), dan stemming (NLTK dan Sastrawi). Hasilnya adalah data yang bersih dan siap dianalisis, seperti ditunjukkan pada Gambar 3.

	content	score	text_clean
6582	Premium nya di gratisin dong nanti kasih bintang...	4	Premium nya di gratisin dong nanti kasih bintang
3059	Sangat ngebas dan banyak pari yasi lagunya sep...	5	Sangat ngebas dan banyak pari yasi lagunya sep...
8843	Makin parah di update malah item iklan ga munc...	2	Makin parah di update malah item iklan ga munc...
268	Seru banget aplikasi paling menyenangkan semua...	4	Seru banget aplikasi paling menyenangkan semua...
1275	Lumayan lama lama gak bisa ulang lagu mane	5	Lumayan lama lama gak bisa ulang lagu mane

[1] Hasil Text Cleaning

	content	score	text_clean	text_lower
6582	Premium nya di gratisin dong nanti kasih bintang...	4	Premium nya di gratisin dong nanti kasih bintang	premium nya di gratisin dong nanti kasih bintang
3059	Sangat ngebas dan banyak pari yasi lagunya sep...	5	Sangat ngebas dan banyak pari yasi lagunya sep...	sangat ngebas dan banyak pari yasi lagunya sep...
8843	Makin parah di update malah item iklan ga munc...	2	Makin parah di update malah item iklan ga munc...	makin parah di update malah item iklan ga munc...
268	Seru banget aplikasi paling menyenangkan semua...	4	Seru banget aplikasi paling menyenangkan semua...	seru banget aplikasi paling menyenangkan semua...
1275	Lumayan lama lama gak bisa ulang lagu mane	5	Lumayan lama lama gak bisa ulang lagu mane	lumayan lama lama gak bisa ulang lagu mane

[2] Hasil Lower Case

	content	score	text_clean	text_lower	tokens_text
6582	Premium nya di gratisin dong nanti kasih bintang...	4	Premium nya di gratisin dong nanti kasih bintang	premium nya di gratisin dong nanti kasih bintang	[premium, nya, di, gratisin, dong, nanti, kasih, bintang]
3059	Sangat ngebas dan banyak pari yasi lagunya sep...	5	Sangat ngebas dan banyak pari yasi lagunya sep...	sangat ngebas dan banyak pari yasi lagunya sep...	[sangat, ngebas, dan, banyak, pari, yasi, lagunya, sep, ...]
8843	Makin parah di update malah item iklan ga munc...	2	Makin parah di update malah item iklan ga munc...	makin parah di update malah item iklan ga munc...	[makin, parah, di, update, malah, item, iklan, ga, munc, ...]
268	Seru banget aplikasi paling menyenangkan semua...	4	Seru banget aplikasi paling menyenangkan semua...	seru banget aplikasi paling menyenangkan semua...	[seru, banget, aplikasi, paling, menyenangkan, semua, ...]
1275	Lumayan lama lama gak bisa ulang lagu mane	5	Lumayan lama lama gak bisa ulang lagu mane	lumayan lama lama gak bisa ulang lagu mane	[lumayan, lama, lama, gak, bisa, ulang, lagu, mane, ...]

[3] Hasil Tokenisasi

	content	score	text_clean	text_lower	tokens_text	stopwords_text
6582	Premium nya di gratisin dong nanti kasih bintang...	4	Premium nya di gratisin dong nanti kasih bintang	premium nya di gratisin dong nanti kasih bintang	[premium, nya, di, gratisin, dong, nanti, kasih, bintang]	[premium, nya, gratisin, kasih, bintang]
3059	Sangat ngebas dan banyak pari yasi lagunya sep...	5	Sangat ngebas dan banyak pari yasi lagunya sep...	sangat ngebas dan banyak pari yasi lagunya sep...	[sangat, ngebas, dan, banyak, pari, yasi, lagunya, sep, ...]	[ngebas, pari, yasi, lagunya, goyang, nasi, pa, ...]
8843	Makin parah di update malah item iklan ga munc...	2	Makin parah di update malah item iklan ga munc...	makin parah di update malah item iklan ga munc...	[makin, parah, di, update, malah, item, iklan, ga, munc, ...]	[parah, update, item, iklan, ga, muncul, kama, ...]
268	Seru banget aplikasi paling menyenangkan semua...	4	Seru banget aplikasi paling menyenangkan semua...	seru banget aplikasi paling menyenangkan semua...	[seru, banget, aplikasi, paling, menyenangkan, semua, ...]	[seru, banget, aplikasi, menyenangkan, semua, ...]
1275	Lumayan lama lama gak bisa ulang lagu mane	5	Lumayan lama lama gak bisa ulang lagu mane	lumayan lama lama gak bisa ulang lagu mane	[lumayan, lama, lama, gak, bisa, ulang, lagu, mane, ...]	[lumayan, gak, ulang, lagu, mane]

[4] Hasil Stopword

	content	score	text_clean	text_lower	tokens_text	stopwords_text	text_stemming
6582	Premium nya di gratisin dong nanti kasih bintang...	4	Premium nya di gratisin dong nanti kasih bintang	premium nya di gratisin dong nanti kasih bintang	[premium, nya, di, gratisin, dong, nanti, kasih, bintang]	[premium, nya, gratisin, kasih, bintang]	[premium, nya, gratisin, kasih, bintang]
3059	Sangat ngebas dan banyak pari yasi lagunya sep...	5	Sangat ngebas dan banyak pari yasi lagunya sep...	sangat ngebas dan banyak pari yasi lagunya sep...	[sangat, ngebas, dan, banyak, pari, yasi, lagunya, sep, ...]	[ngebas, pari, yasi, lagunya, goyang, nasi, pa, ...]	[ngebas, pari, yasi, lagu, goyang, nasi, padang]
8843	Makin parah di update malah item iklan ga munc...	2	Makin parah di update malah item iklan ga munc...	makin parah di update malah item iklan ga munc...	[makin, parah, di, update, malah, item, iklan, ga, munc, ...]	[parah, update, item, iklan, ga, muncul, kama, ...]	[parah, update, item, iklan, ga, muncul, kama, ...]
268	Seru banget aplikasi paling menyenangkan semua...	4	Seru banget aplikasi paling menyenangkan semua...	seru banget aplikasi paling menyenangkan semua...	[seru, banget, aplikasi, paling, menyenangkan, semua, ...]	[seru, banget, aplikasi, menyenangkan, semua, ...]	[seru, banget, aplikasi, senang, musik, lagu, ...]
1275	Lumayan lama lama gak bisa ulang lagu mane	5	Lumayan lama lama gak bisa ulang lagu mane	lumayan lama lama gak bisa ulang lagu mane	[lumayan, lama, lama, gak, bisa, ulang, lagu, mane, ...]	[lumayan, gak, ulang, lagu, mane]	[lumayan, gak, ulang, lagu, mane]

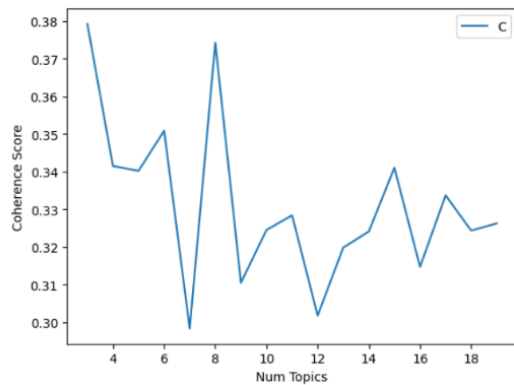
[5] Hasil Stemming

Gambar 3. Hasil Proses Text Preprocessing

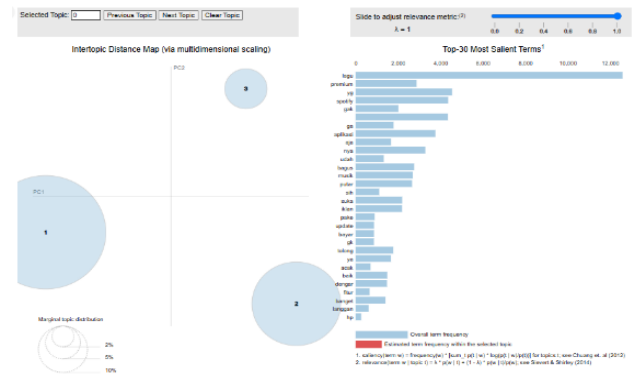
Pada Gambar 3 menampilkan transformasi hasl data teks pada tahapan text preprocessing, tahapan ini mencakup pembersihan teks dari karakter khusus, penghapusan tanda baca dan angka, normalisasi huruf menjadi huruf kecil, tokenisasi kata, penghapusan stopwords, dan stemming. Pada proses ini penting untuk mengurangi noise dalam teks dan meningkatkan kualitas data.

3.3. Hasil Metode *Latent Dirichlet Allocation* (LDA)

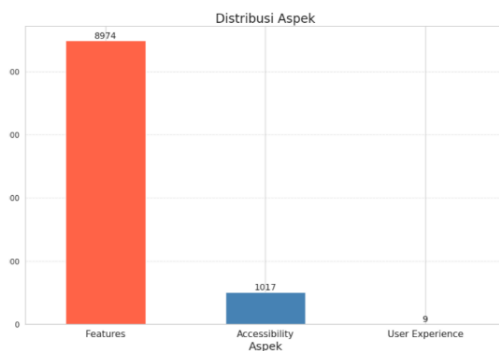
Setelah preprocessing, data diproses dengan *Latent Dirichlet Allocation* (LDA) untuk menentukan num topics dan topik dominan beserta aspeknya. Tahapan ini mencakup pembuatan model dan kamus LDA, perhitungan skor koherensi, visualisasi pyLDavis, identifikasi *keywords*, dan penentusn aspek. Hasilnya diperoleh 3 aspek utama yaitu Accessibility, Features, dan *User Experience*. Hasil dari tahapan proses metode *Latent Dirichlet Allocation* (LDA) dapat dilihat pada Gambar 4 berikut.



[6] Visualisasi Koherensi



[7] Visualisasi pyLDAvis



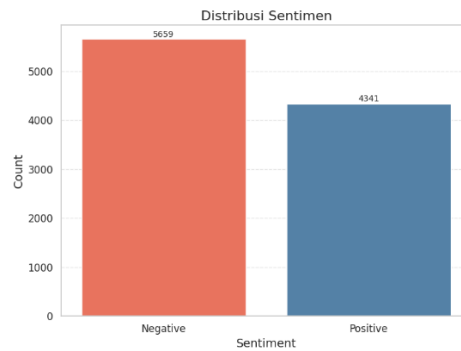
[8] Visualisasi Jumlah Aspek

Gambar 4. Hasil Metode *Latent Dirichlet Allocation*

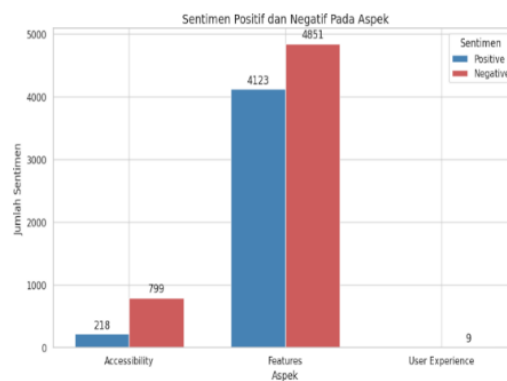
Gambar 4 menyajikan hasil visualisasi topic modelling dengan menggunakan metode *Latent Dirichlet Allocation* (LDA) yang menunjukkan bahwa pemodelan dengan 3 topik menghasilkan nilai koherensi tertinggi. Nilai koherensi digunakan sebagai metrik untuk mengukur kualitas topik yang dihasilkan, semakin tinggi nilai koherensi akan semakin konsisten dan bermakna untuk topik tersebut. Visualisasi dengan menggunakan pyLDAvis menunjukkan distribusi topik dalam bentuk lingkaran dan kata kunci dominan pada tiap topik. Interpretasi terhadap kata – kata dominan dalam masing – masing topik menunjukkan bahwa topik pertama berkaitan dengan fitur aplikasi (*Features*), topik kedua berkaitan dengan aksesibilitas (*Accessibility*), dan topik ketiga yang berhubungan dengan pengalaman pengguna secara umum (*User Experience*). Pemisahan ini digunakan untuk mendefinisikan aspek – aspek utama dalam analisis sentiment berbasis aspek (ABSA).

3.4. Hasil Labelling Data

Proses pelabelan data menggunakan metode *InSet Lexicon*. Metode ini merupakan kumpulan kata secara khusus dirancang untuk analisis sentiment berbahasa Indonesia. Ada dua hasil dari pelabelan data dengan *InSet Lexicon* yaitu sentimen positif dan sentiment negative, dengan nilai score polaritas lebih dari 0 masuk ke dalam sentiment positif sedangkan score dengan nilai di bawah 0 atau minus masuk ke dalam sentiment negative. Kemudian dicari nilai sentimen untuk setiap aspeknya. Berikut hasil pelabelan data.



[9] Visualisasi Sentimen



[10] Visualisasi Jumlah Sentimen Setiap Aspek



[11] Wordcloud Sentimen Positif



[12] Wordcloud Sentimen Negatif

Gambar 5. Hasil Labelling Data

Pada Gambar 5 menampilkan hasil dari proses dari Labelling Data. Visualisasi sentimen menghasilkan sentimen positif 4.341 ulasan, dan sentimen negatif 5.659 ulasan, kemudian untuk hasil sentimen setiap aspek menghasilkan jumlah sentimen untuk 3 aspek antara lain aspek Accessibility dengan jumlah 218 untuk positif dan 799 untuk negatif, aspek Features 4.123 untuk positif dan 4851 untuk negatif, kemudian yang terakhir aspek User Experience dengan jumlah hanya mendapat 9 untuk negatif, ini dikarenakan jumlah data yang sedikit sehingga tidak bisa optimal. Wordcloud yang menggambarkan kata – kata yang paling sering muncul dalam ulasan pengguna berdasarkan klasifikasi sentimen. Pada wordcloud sentimen

positif terdapat kata – kata seperti “bagus”, “suka”, dan “gratis”, yang menunjukkan bahwa banyak pengguna merasa puas terhadap aplikasi Spotify. Sementara itu wordcloud sentimen negative terdapat kata – kata seperti “iklan”, “tolong”, “bayar” ini dapat menunjukkan keluhan pengguna terhadap aplikasi. Meskipun pendekatan leksikon ini praktis dan cukup efisien, perlu dicatat bahwa proses pelabelan ini tidak disertai dengan validasi manual oleh annotator, keterbatasan ini muncul karena keterbatasan waktu dalam proses penelitian. Hal ini dapat berimplikasi pada ketidaktepatan pelabelan, terutama terhadap ulasan yang mengandung konteks ganda, ironi, atau sarkasme yang tidak dapat dikenali oleh pendekatan leksikal sederhana. Oleh karena itu perlu dilakukannya validasi manual terhadap sampel ulasan atau menggabungkan pendekatan berbasis konteks untuk meningkatkan akurasi dan ketepatan klasifikasi sentiment.

3.5. Hasil Evaluasi Data

Dalam proses evaluasi data menggunakan klasifikasi Random Forest, ada 2 tahap yaitu evaluasi untuk sentimen dan evaluasi untuk aspek. Evaluasi sentimen akan dilakukan proses pembobotan TF-IDF, yang kemudian akan dibagi menjadi 80% data latih dan 20% data uji, dengan oversampling SMOTE. Untuk evaluasi aspek sendiri memiliki metode klasifikasi yang sama dengan evaluasi sentiment dan pembagian data yang sama dengan SMOTE. Hasil nya akan menunjukkan performa dari klasifikasi dalam menganalisis setiap parameter – parameternya. Hasil evaluasi data dapat di lihat di bawah ini.

Tabel 1. Hasil Evaluasi Sentimen Dengan Random Forest

Sentiment	Precision	Recall	F1-Score	Accuracy
Negative	0.94	0.96	0.95	0.94
Positive	0.95	0.92	0.93	0.94

Tabel 2. Hasil Evaluasi Aspek Dengan Random Forest

Topik	Aspect	Accuracy	Precision	Recall	F1-score
0	Features	0.948747	0.948803	0.948747	0.948764
1	Accessibility	0.882353	0.885291	0.882353	0.867391
2	User Experience	1.000000	1.000000	1.000000	1.000000

3.6. Hasil Parameter Terbaik

Untuk mengevaluasi kinerja model klasifikasi sentimen secara lebih menyeluruh dan menghindari overfitting, penelitian ini menggunakan Teknik K-Fold Cross Validation dengan jumlah fold sebanyak 5. Metode ini bertujuan untuk memastikan bahwa model memiliki performa yang stabil dan dapat digeneralisasi dengan baik terhadap data yang belum pernah dilihat sebelumnya. Dalam implementasinya dataset dibagi menjadi 5 lipatan (fold), Dimana pada setiap iterasi 4 fold digunakan sebagai data latih dan 1 fold digunakan sebagai data uji. Proses ini dilakukan sebanyak 5 kali sehingga setiap fold akan berperan sebagai data uji satu kali. Validasi silang ini dilakukan untuk evaluasi performa algoritma Random Forest dalam klasifikasi sentimen. Adapun parameter yang digunakan pada Random Forest yaitu $n_estimators = 100$, $random_state = 42$. Dalam setiap iterasi dilakukan juga oversampling dengan menggunakan teknik SMOTE pada data latih untuk mengatasi ketidakseimbangan kelas antara sentiment positif dan negatif. Kemudian model dilatih menggunakan data hasil SMOTE dan diuji pada data uji asli. Metrik evaluasi yang digunakan meliputi Accuracy, Precision, Recall, dan F1-score, dengan pendekatan rata – rata berbobot untuk

memperhitungkan distribusi kelas yang tidak seimbang. Hasil dapat dilihat pada Tabel 3 berikut.

Tabel 3. Hasil Fold Dengan K-Fold Cross Validation

Fold	Accuracy	Precision	Recall	F1-score
1	0.934105	0.934750	0.934105	0.933842
2	0.945171	0.945145	0.945171	0.945129
3	0.939135	0.939099	0.939135	0.939088
4	0.927529	0.927572	0.927529	0.927388
5	0.942124	0.942091	0.942124	0.942096

Hasil K-Fold Cross Validation diatas menunjukkan nilai setiap parameter dari 5 fold. Untuk fold terbaik berdasarkan F1-score diantara ke 5 fold didapatkan oleh fold 2 dengan jumlah F1-score tertinggi 94.51%, model ini menunjukkan kinerja performa yang sangat baik pada fold 2 dan keseimbangan yang bagus antara Precision dan Recall.

4. Kesimpulan

Penelitian ini berhasil mengimplementasikan metode Latent Dirichlet Allocation (LDA) dalam analisis sentimen berbasis aspek (Aspect-Based Sentiment Analysis) terhadap ulasan pengguna aplikasi Spotify. Berdasarkan hasil ekstraksi topik menggunakan LDA, diperoleh tiga aspek utama yang paling sering muncul dalam ulasan yaitu Features, Accessibility, dan User Experience. Kemudian untuk pelabelan sentimen dengan menggunakan kamus InSet Lexicon menunjukkan bahwa dari total 10.000 data ulasan terdapat 4.341 ulasan dengan sentimen positif dan 5.659 ulasan dengan sentimen negatif. Hal ini menunjukkan bahwa secara umum ulasan pengguna cenderung lebih banyak mengandung kritik atau keluhan terhadap aplikasi Spotify. Lalu untuk evaluasi model klasifikasi menggunakan algoritma Random Forest menghasilkan akurasi tinggi baik dalam klasifikasi sentimen maupun aspek, dengan akurasi keseluruhan mencapai 94% untuk hasil sentimen. Validasi menggunakan K-Fold Cross Validation menunjukkan bahwa model memiliki performa yang stabil dan tidak overfitting dengan fold ke-2 menghasilkan parameter terbaik. Penelitian ini berkontribusi dalam memberikan pendekatan terstruktur untuk memahami persepsi pengguna berdasarkan aspek tertentu dari ulasan. Hasil analisis dapat digunakan oleh pengembang Spotify untuk mengidentifikasi area layanan yang perlu diperbaiki secara spesifik. Namun dalam penelitian ini masih terdapat beberapa kekurangan, diantaranya adalah pelabelan sentimen yang sepenuhnya dilakukan secara otomatis menggunakan kamus InSet Lexicon tanpa melalui proses validasi manual oleh annotator. Hal ini berpotensi menyebabkan kesalahan klasifikasi, terutama pada ulasan yang mengandung ambigu, atau konteks ganda. Kemudian untuk hasil pada aspek User Experience kurang valid dikarenakan kurang nya data, sehingga membuat hasil kurang optimal. Oleh karena itu, penelitian selanjutnya disarankan untuk melibatkan proses validasi manual ataupun mengadopsi pendekatan klasifikasi berbasis konteks seperti model transformer, misalnya IndoBERT ini berguna untuk meningkatkan ketepatan pelabelan sentimen, dan menambahkan data ulasan jika ada kelas yang memiliki data sedikit atau kekurangan data .

Referensi

- Aritonang, P. A., Johan, M. E., & Prasetiawan, I. (2022). Aspect-Based Sentiment Analysis on Application Review using Convolutional Neural Network. *Ultima InfoSys : Jurnal Ilmu Sistem Informasi*, 13(1), 54–61. <https://doi.org/10.31937/si.v13i1.2684>
- Darusalam, T., Alam, S., Komara, M. A., Informatika, T., Tinggi, S., Wastukencana, T.,

- Purwakarta, K., Barat, J., Sentimen, A., & Gain, I. (2024). *ANALISIS SENTIMEN PENGGUNA SLIK OJK MENGGUNAKAN NAÏVE BAYES*. 8(5), 8709–8716.
- Fauzi, F., Abdulloh, F. F., Pambudi, I. R., Yogyakarta, U. A., Sleman, K., Informatika, J., Ilmu, F., Universitas, K., Yogyakarta, A., & Machine, S. V. (2021). *ANALISIS SENTIMEN PENGGUNA YOUTUBE TERHADAP PROGRAM VAKSIN COVID-19*. 141–148.
- Fauziah, A. N. (2023). Analisis sentimen menggunakan naïve bayes classifier dan inset lexicon pada twitter (studi kasus: Mie Gacoan). *Repository.Uinjkt.Ac.Id*. [https://repository.uinjkt.ac.id/dspace/handle/123456789/76248%0Ahttps://repository.uinjkt.ac.id/dspace/bitstream/123456789/76248/1/Aisyah Nur Fauziah-FST.pdf](https://repository.uinjkt.ac.id/dspace/handle/123456789/76248%0Ahttps://repository.uinjkt.ac.id/dspace/bitstream/123456789/76248/1/Aisyah%20Nur%20Fauziah-FST.pdf)
- Hikmah, F. N., Basuki, S., & Azhar, Y. (2020). Deteksi Topik Tentang Tokoh Publik Politik Menggunakan Latent Dirichlet Allocation (LDA). *Jurnal Repositor*, 2(4), 415–426. <https://doi.org/10.22219/repositor.v2i4.52>
- Karyono, Z. R., Mursityo, Y. T., & Az-Zahra, H. M. (2019). Analisis Perbandingan Pengalaman Pengguna Pada Aplikasi Music Streaming Menggunakan Metode UX Curve (Studi Pada Spotify dan JOOX). *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 3(7), 6422–6429. <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/5721>
- Khairunnisa, S., Adiwijaya, A., & Faraby, S. Al. (2021). Pengaruh Text Preprocessing terhadap Analisis Sentimen Komentar Masyarakat pada Media Sosial Twitter (Studi Kasus Pandemi COVID-19). *Jurnal Media Informatika Budidarma*, 5(2), 406. <https://doi.org/10.30865/mib.v5i2.2835>
- Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. (2022). *Analisis Sentimen Ulasan Aplikasi Dana dengan Metode Random Forest*. 6(9), 4305–4313.
- Mustofa, A., & Novita, R. (2022). Klasifikasi Sentimen Masyarakat Terhadap Pemberlakuan Pembatasan Kegiatan Masyarakat Menggunakan Text Mining Pada Twitter. *Building of Informatics, Technology and Science (BITS)*, 4(1), 200–208. <https://doi.org/10.47065/bits.v4i1.1628>
- Purnomo, A. (2022). *Impementasi Web Scraping Pada OJS Dengan Metode CSS Selector*. 3(2), 63–68.
- Rahma Yustihan, S., & Pandu Adikara, P. (2021). Analisis Sentimen berbasis Aspek terhadap Data Ulasan Rumah Makan menggunakan Metode Support Vector Machine (SVM). *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 5(3), 1017–1023. <http://j-ptiik.ub.ac.id>
- Sidiq, S., Maburur, N. S., Informatika, S. T., Teknik, F., & Tangerang, U. M. (2025). *Pengembangan Model Prediksi Risiko Diabetes Menggunakan Pendekatan AdaBoost dan Teknik Oversampling*. 4, 13–23.
- Yulita, W. (2021). Analisis Sentimen Terhadap Opini Masyarakat Tentang Vaksin Covid-19 Menggunakan Algoritma Naïve Bayes Classifier. *Jurnal Data Mining Dan Sistem Informasi*, 2(2), 1. <https://doi.org/10.33365/jdmsi.v2i2.1344>