

ANALISIS SENTIMEN BERBASIS ASPEK PADA SISTEM LAYANAN PENGADUAN MASYARAKAT DI KOTA SURABAYA MENGGUNAKAN METODE LATENT DIRICHLET ALLOCATION DAN NAÏVE BAYES

Ferdy Atmaja, Agussalim, Eka Dyar Wahyuni

Sistem Informasi, Universitas Pembangunan Nasional "Veteran" Jawa Timur

Jalan Raya Rungkut Madya, No. 1, Gunung Anyar, Surabaya, Indonesia

Ferdy.atmaja308@gmail.com

ABSTRAK

Sistem layanan pengaduan masyarakat menjadi salah satu sarana penting bagi pemerintah kota untuk menerima, memproses, dan menindaklanjuti berbagai keluhan warga. Namun, peningkatan volume data pengaduan di Kota Surabaya menjadi tantangan dalam pengelolaan manual yang kurang efektif. Analisis mendalam terhadap data pengaduan diperlukan untuk memahami opini masyarakat terhadap kualitas layanan. Penelitian ini bertujuan untuk melakukan analisis sentimen berbasis aspek terhadap data aduan. Penelitian dilakukan dengan memanfaatkan metode *Latent Dirichlet Allocation* (LDA) untuk mengidentifikasi aspek utama keluhan dan menggunakan *Multinomial Naïve Bayes* (MNB) untuk mengklasifikasikan sentimen. Data yang digunakan terdiri dari 10.847 pengaduan yang dikumpulkan melalui aplikasi WargaKu dan Media Center sepanjang tahun 2023. Metode LDA berhasil mengidentifikasi 17 topik utama, termasuk administrasi, infrastruktur, dan informasi publik. Selanjutnya, Model MNB mencapai akurasi 80% setelah penerapan resampling, membuktikan keandalannya dalam klasifikasi sentimen. Hasil penelitian memberikan wawasan yang relevan untuk membantu pemerintah kota memprioritaskan perbaikan pada isu utama serta menawarkan pendekatan yang dapat diadopsi oleh sistem pengaduan serupa.

Kata kunci : Analisis Sentimen, Latent Dirichlet Allocation, Naïve Bayes, Pengaduan Masyarakat, Layanan Public

1. PENDAHULUAN

Pelayanan publik merupakan elemen esensial dalam menjalankan fungsi pemerintahan yang bertujuan untuk memenuhi kebutuhan masyarakat sekaligus menciptakan harmoni antara hak dan kewajiban warga negara. Di era digital, peningkatan kualitas layanan publik menjadi semakin relevan dengan adopsi teknologi informasi untuk mendukung efisiensi, transparansi, dan aksesibilitas layanan [1].

Kota Surabaya telah mengimplementasikan langkah progresif melalui dua platform digital, yaitu Media Center dan aplikasi WargaKu, yang dirancang untuk memfasilitasi masyarakat dalam menyampaikan aspirasi, keluhan, dan masukan terhadap pelayanan publik [2].

Langkah ini tidak hanya mencerminkan komitmen pemerintah untuk mendekatkan diri dengan masyarakat tetapi juga menjawab tantangan dalam mewujudkan good governance berbasis teknologi.

Sejak diperkenalkan, aplikasi WargaKu dan Media Center telah mengalami peningkatan signifikan dalam jumlah aduan yang diterima. Menurut data dari Diskominfo Surabaya, jumlah laporan yang masuk pada tahun pertama mencapai 698, meningkat tajam menjadi 2.717 pada tahun 2012, kemudian 4.176 pada tahun 2013, dan mencapai puncaknya pada tahun 2014 dengan 4.298 keluhan [3].

Hal ini menunjukkan bahwa kedua platform ini menjadi sumber informasi yang berharga dalam memahami kebutuhan dan ekspektasi masyarakat. Namun, pengelolaan data aduan yang terus meningkat

secara manual masih menjadi tantangan yang membutuhkan solusi inovatif. Dalam konteks ini, analisis sentimen menjadi pendekatan yang potensial untuk mengelompokkan dan memahami respons masyarakat secara lebih efektif.

Analisis sentimen memungkinkan pemerintah untuk mengidentifikasi pola, tren, dan sentimen masyarakat terhadap layanan publik yang diberikan. Dengan mengintegrasikan metode klasifikasi seperti *Multinomial Naïve Bayes* (MNB) dan metode analisis teks menggunakan *Latent Dirichlet Allocation* (LDA), penelitian ini bertujuan untuk memberikan wawasan mendalam mengenai persepsi masyarakat terhadap berbagai aspek layanan. Pendekatan ini tidak hanya membantu pemerintah untuk merespons kebutuhan masyarakat secara tepat waktu tetapi juga memberikan kontribusi pada peningkatan kualitas layanan publik yang berbasis data.

Penelitian ini difokuskan pada pengolahan data aduan masyarakat di Surabaya yang dikumpulkan melalui aplikasi WargaKu dan Media Center sepanjang tahun 2023. Data ini akan dianalisis untuk mengidentifikasi aspek utama keluhan melalui LDA serta klasifikasi sentimen untuk memahami pandangan masyarakat. Hasil dari analisis ini diharapkan dapat memberikan rekomendasi perbaikan yang relevan bagi pemerintah dalam upaya meningkatkan kualitas layanan publik serta menjadi referensi bagi pengembangan penelitian serupa di masa depan.

2. TINJAUAN PUSTAKA

Pada tahapan ini, akan dipaparkan beberapa literatur ilmiah yang relevan dengan penelitian.

2.1. Penelitian Terdahulu

Penelitian yang dilakukan oleh Roiqoh dan Zaman (2023) berfokus pada analisis sentimen berbasis aspek terhadap ulasan aplikasi Mobile JKN yang tersedia di *Google Play Store*. Penelitian ini bertujuan untuk memahami penilaian publik terhadap aplikasi tersebut dari berbagai sudut pandang guna memberikan wawasan dalam pengembangan aplikasi di masa depan. Data yang digunakan terdiri dari 10.752 ulasan dari dua versi aplikasi terbaru. Metode analisis melibatkan *Naïve Bayes* untuk klasifikasi sentimen dan *Lexicon Based* (Inset Lexicon) sebagai metode alternatif. Selain itu, pemodelan topik dilakukan untuk menemukan topik tersembunyi dalam ulasan. Hasil penelitian menunjukkan tiga topik utama, yaitu Pelayanan dan Fitur, Register dan Login, serta Kepuasan Pengguna. Dari segi performa, *Naïve Bayes* menunjukkan akurasi tertinggi sebesar 94,75%, mengungguli *Lexicon Based* yang hanya mencapai 59,99% [4].

Penelitian selanjutnya oleh Yutika et al. (2021) menganalisis sentimen pada ulasan produk dari platform Female Daily yang bersifat multilingual. Tujuan penelitian ini adalah untuk membandingkan efektivitas berbagai teknik *preprocessing*, seperti penerjemahan kata asing, penghapusan stopword, dan *stemming*. Data yang digunakan terdiri dari 5.054 ulasan produk kecantikan, meliputi kategori toner, serum, sunscreen, scrub, dan exfoliator. Analisis dilakukan menggunakan pembobotan TF-IDF dan algoritma *Naïve Bayes* dengan evaluasi performa berdasarkan akurasi, presisi, recall, dan F1-Score. Hasil penelitian menunjukkan bahwa penerjemahan ulasan dari Bahasa Inggris ke Bahasa Indonesia meningkatkan performa analisis sentimen. Selain itu, tidak menggunakan penghapusan stopword menghasilkan F1-Score terbaik sebesar 62,81% [5].

2.2. Media Center & Aplikasi Wargaku

Media center merupakan sistem layanan terpadu yang diperuntukkan untuk warga Surabaya yang ingin ikut serta dalam kemajuan pembangunan kota Surabaya. Melalui media center, masyarakat dapat memantau progress pembangunan berjalan sejauh mana dan memastikan bahwa pelaksanaan sudah sesuai dengan target pemerintah [6].

WargaKu adalah akronim dari Wadah Aspirasi Rukun Warga Rukun Tetangga dan Kampung Unggul. Aplikasi WargaKu merupakan bagian dari media center yang dikembangkan dalam aplikasi mobile. Melalui aplikasi WargaKu, masyarakat dapat menyampaikan pengaduan yang terbagi dalam 5 kategori yaitu apresiasi, keluhan, kritik, permohonan informasi, dan saran [7].

Dalam penyampaian pengaduan melalui Aplikasi WargaKu, masyarakat dapat memilih topik yang akan

dihubungkan Organisasi Perangkat Daerah (OPD) terkait, serta melampirkan informasi lengkap mengenai aduan, termasuk lokasi dan lampiran yang diperlukan untuk mempermudah dalam proses tindak lanjut. OPD terkait akan memberikan tanggapan dalam kurun waktu kurang dari 24 jam sesuai dengan pengaduan yang diajukan, dan setelahnya, masyarakat dan OPD dapat berinteraksi secara langsung melalui aplikasi WargaKu.

2.3. Analisis Sentimen Berbasis Aspek

Aspect-Based Sentiment Analysis (ABSA), atau yang dikenal dalam bahasa Indonesia sebagai Analisis Sentimen Berbasis Aspek, merupakan cabang dari analisis sentimen yang memungkinkan penguraian opini secara lebih mendalam terhadap aspek-aspek tertentu dalam sebuah dokumen [8].

Pendekatan ini bertujuan menganalisis sentimen yang diarahkan pada aspek spesifik dari suatu entitas, seperti opini pengguna atau, dalam penelitian ini, aduan terkait aspek tertentu. ABSA sering diterapkan dalam domain spesifik seperti ulasan produk, layanan publik, atau media sosial untuk menggali wawasan yang lebih kaya dan bermakna dibandingkan sekadar memahami sentimen secara umum [9].

2.4. Term Frequency–Inverse Document Frequency (TF–IDF)

TF-IDF adalah metode pembobotan yang umum digunakan untuk mengubah data teks menjadi format numerik. Teknik ini menggabungkan dua algoritma pembobotan, yaitu TF yang menghitung frekuensi kemunculan kata dalam sebuah dokumen, dan IDF yang mengukur seberapa penting kata tersebut dalam seluruh kumpulan dokumen [10].

2.5. Multinomial Naïve Bayes

Multinomial Naive Bayes adalah teknik yang sering digunakan untuk mengklasifikasikan teks, terutama dalam pengolahan bahasa alami (NLP). Metode ini bekerja dengan menghitung seberapa sering sebuah kata muncul dalam dokumen dan menggunakan data tersebut untuk memprediksi kelas dokumen. Dalam prosesnya, metode ini menganggap semua kata dalam dokumen bergantung pada kelas yang dianalisis, tetapi tidak memperhatikan hubungan antar kata. Multinomial Naive Bayes memanfaatkan tidak hanya keberadaan kata, tetapi juga seberapa sering kata tersebut muncul, sehingga sangat cocok untuk mengklasifikasikan teks dan dokumen berdasarkan isinya [11].

2.6. Topic Modeling

Topic modelling merupakan metode analisis teks yang bertujuan untuk mengelompokkan dokumen berdasarkan kemiripan topik [12].

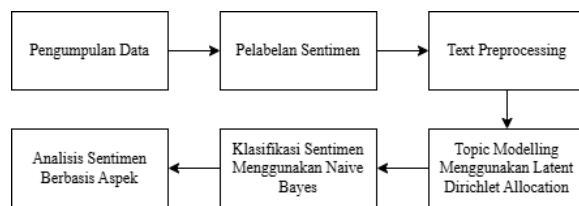
Teknik ini membantu menemukan topik tersembunyi dan kata-kata yang mewakili topik dalam kumpulan dokumen. Setiap dokumen dianggap sebagai kombinasi dari beberapa topik dengan

proporsi tertentu, sedangkan setiap topik terdiri dari sekumpulan kata yang memiliki hubungan erat dengan topik tersebut. Hasil akhirnya adalah kumpulan topik yang sering muncul dalam dokumen berdasarkan pola yang teridentifikasi.

Latent Dirichlet Allocation (LDA) adalah salah satu metode populer dalam topic modelling. LDA bekerja dengan prinsip bahwa setiap dokumen adalah campuran acak dari berbagai topik, di mana setiap topik memiliki distribusi kata unik [13]. Dengan pendekatan probabilistik, LDA menggunakan distribusi Dirichlet untuk mengalokasikan kata-kata dalam dokumen ke berbagai topik.

3. METODE PENELITIAN

Pada penelitian ini akan dilakukan beberapa tahapan diantaranya melakukan pengumpulan data, pelabelan sentimen, *text preprocessing* yang meliputi *case folding*, *cleansing*, *normalization*, *stopword removal*, *stemming*, dan *tokenizing*. Kemudian dilanjutkan dengan tahapan topic modeling, klasifikasi sentiment, dan dilanjutkan dengan proses analisis sentiment berbasis aspek. Adapun alur penelitian ini tersaji pada Gambar 1 berikut.



Gambar 1. Alur penelitian

3.1. Pengumpulan Data

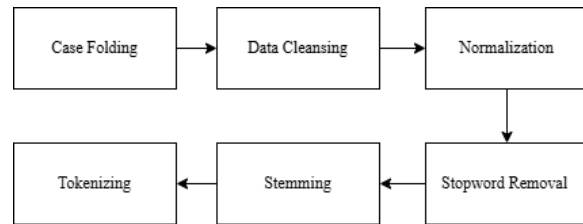
Data didapat dari sumber resmi yaitu PPID Surabaya, yang merupakan instansi penyelenggara layanan pengaduan masyarakat. Informasi yang dikumpulkan meliputi isi laporan, tanggal pelaporan, kategori laporan, topik laporan, dan organisasi perangkat daerah (OPD) terkait.

3.2. Pelabelan Sentimen

Pada tahap ini akan dilakukan pelabelan sentimen yang dilakukan secara manual oleh 3 pelabel yang masing-masing melakukan labeling untuk seluruh baris data

3.3. Text Preprocessing

Tahapan text preprocessing melibatkan beberapa tahapan diantaranya adalah: *Case folding*, *data cleansing*, *normalization*, *stopwords removal*, *stemming*, dan *tokenizing*. Alur *preprocessing* digambarkan pada Gambar 2 berikut:



Gambar 2. Alur Text Preprocessing

3.4. Topic Modeling

Pada tahap ini, metode *Latent Dirichlet Allocation* (LDA) digunakan untuk mengungkap pola-pola topik yang tersembunyi dalam keluhan warga Surabaya, dengan tujuan untuk mengidentifikasi aspek-aspek utama yang menjadi perhatian masyarakat.

3.5. Klasifikasi Sentimen

Pada tahap ini, metode *Naïve Bayes* diimplementasikan dengan tujuan mengklasifikasikan sentimen dalam keluhan warga menjadi kategori tertentu berdasarkan pola kata dalam teks.

3.6. Analisis Sentimen Berbasis Aspek

Pada tahap ini, sentimen dalam setiap keluhan dianalisis secara terfokus untuk tiap aspek yang telah diinterpretasikan sebelumnya. Analisis sentimen pada tahap terakhir dapat memberikan wawasan spesifik tentang persepsi warga terhadap setiap aspek, memungkinkan identifikasi isu yang paling sering dikeluhkan dan aspek dengan sentimen negatif tinggi, yang dapat menjadi dasar rekomendasi peningkatan layanan atau kebijakan terkait.

4. HASIL DAN PEMBAHASAN

Penelitian ini dilakukan menggunakan *Google Colaboratory* dengan bahasa pemrograman Python. Untuk analisis sentimen, digunakan algoritma *Multinomial Naïve Bayes* (MNB) untuk membangun model klasifikasi sentimen. Selain itu, metode *Latent Dirichlet Allocation* (LDA) digunakan untuk menemukan topik utama dari data keluhan. Proses dan hasil dari kedua pendekatan ini diuraikan lebih lanjut pada bagian berikut.

4.1. Pengumpulan Data

Data yang digunakan dalam penelitian ini merupakan data aduan masyarakat yang disampaikan melalui Aplikasi Wargaku dan Media Center Kota Surabaya. Data diperoleh secara resmi dari Pejabat Pengelola Informasi dan Dokumentasi (PPID) Kota Surabaya.

Kategori	Tanggal Keluhan	Keluhan	Instansi
Keluhan	2023-01-01 08:26:00	tidak pernah mendapat bantuan apa2. suami dan saya terkena	Dinas Sosial
Keluhan	2023-01-01 08:36:00	Lapor gorong-gorong saluran dari rt09rw05 depan jln Kalikepiting 33 1. Taman Prestasi tempat wudhu	Dinas Sumber Daya Air dan Bina Marga
Kritik	2023-01-01 08:45:00	Kotor.\n2. Ada genangan di bawah sudah di kerjakan, tapi kok	Dinas Lingkungan Hidup
Keluhan	2023-01-01 09:04:00	sepertinya setengah hati atau	Dinas Sumber Daya Air dan Bina Marga
Keluhan	2023-01-01 09:44:00	saluran jebol	Dinas Sumber Daya Air dan Bina Marga
Keluhan	2023-01-01 12:07:00	pengangkutan 13 titik sampah liar dialamat: \n#Kalibokor no 35, 79,\n#Kalibokor Selatan no 92, 54, 46, 36, 34, 26.\n#Nagel Jaya Barat	Dinas Lingkungan Hidup

Gambar 3. Data aduan masyarakat Kota Surabaya

Gambar 3 menunjukkan data yang diterima dari PPID Kota Surabaya yang kemudian disimpan dalam format file .xlsx. Jumlah total data yang digunakan terdiri dari 10.847 record aduan. Data merupakan data aduan yang masuk selama periode Januari hingga Desember 2023.

4.2. Pelabelan Sentimen

Pelabelan sentimen dilakukan secara manual oleh 3 pelabel yang masing-masing melakukan labeling untuk seluruh baris data. Hasil pelabelan menunjukkan bahwa terdapat 39 keluhan teridentifikasi sebagai positif, 6.208 negatif, 4.445 netral, dan 39 diberi label (-) karena dianggap tidak bermakna atau tidak memberikan informasi yang jelas terkait sentimen. Hasil pelabelan dapat dilihat pada Tabel 1.

Tabel 1. Hasil pelabelan

Kategori	Sebelum
Positif	56
Negatif	6.208
Netral	4.445
-	39
Total	10.748

Mengingat bahwa data yang dianalisis adalah keluhan masyarakat, di mana umumnya orang cenderung melaporkan masalah atau ketidakpuasan ketimbang memberikan pujian, hal ini menyebabkan ketidakseimbangan data. Ketidakseimbangan tersebut berpotensi memengaruhi keakuratan hasil analisis sentimen secara keseluruhan. Oleh karena itu, keluhan yang terlabel positif dan yang diberi label (-) akan dihapus dari dataset. Tabel 2 memperlihatkan jumlah data hasil pelabelan yang sudah dilakukan pengurangan.

Tabel 2. Hasil pengurangan setelah pelabelan

Kategori	Sebelum	Sesudah
Positif	56	0
Negatif	6.208	6.208
Netral	4.445	4.445
-	39	0
Total	10.748	10.653

4.3. Text Preprocessing

Text preprocessing dilakukan dengan bantuan library Python dan melalui beberapa tahapan meliputi *case folding*, *data cleansing*, *normalization*, *stopword removal*, *stemming*, dan *tokenizing*.

4.3.1. Case Folding

Proses *case folding* pada penelitian ini bertujuan untuk mengubah seluruh teks menjadi huruf kecil (*lowercase*). Hasil *case folding* dapat dilihat pada Tabel 3.

Tabel 3. Case Folding

Sebelum	Sesudah
Saya ingin mengajukan KIP	saya ingin mengajukan kip

4.3.2. Data Cleansing

Tahapan ini dilakukan untuk memastikan teks bersih dari elemen-elemen yang tidak relevan. Data cleansing dilakukan secara bertahap mulai dari menghilangkan tag HTML, emoji, white space, tanda baca, serta karakter-karakter lainnya yang tidak diperlukan. Hasil dari proses cleansing dapat dilihat pada Tabel 4.

Tabel 4. Data Cleansing

Sebelum	Sesudah
1. taman prestasi tempat wudhu kotor. 2. ada genangan di bawah ayunan	taman prestasi tempat wudhu kotor ada genangan di bawah ayunan

4.3.3. Normalization

Tahapan ini bertujuan untuk mengubah data teks aduan yang masih belum baku menjadi baku. Hasil dari tahapan normalisasi dapat dilihat pada table 5.

Tabel 5. Normalization

Sebelum	Sesudah
bgmn cara dftr epeken	bagaimana cara daftar epeken

4.3.4. Stopword Removal

Tahapan ini bertujuan untuk mengeliminasi kata-kata yang dianggap tidak penting, seperti kata hubung dan kata umum yang tidak memberikan makna signifikan terhadap konteks analisis. Hasil dari tahapan ini dapat dilihat pada table 6.

Tabel 6. Stopword Removal

Sebelum	Sesudah
apa program berobat gratis dengan ktp surabaya masih berlaku dan apakah bila saya terlambat bayar iuran saya tidak bisa berobat rumah sakit	apa program berobat gratis ktp surabaya berlaku bila terlambat bayar iuran berobat rumah sakit

4.3.5. Stemming

Tahapan *stemming* dilakukan untuk mengubah kata-kata dalam teks ke bentuk dasarnya (*root word*). Hasil dari tahapan ini dapat dilihat pada Tabel 7.

Tabel 7. Stemming

Sebelum	Sesudah
mohon bantuan rekan dlh pengangkutan sampah jalan kalibokor selatan	mohon bantu rekan dlh angkut sampah jalan kalibokor selatan

4.3.6. Tokenizing

Tahapan *tokenizing* digunakan untuk memecah teks menjadi unit-unit yang lebih kecil. Hasil dari proses *tokenizing* dapat dilihat pada Tabel 8.

Tabel 8. *Tokenizing*

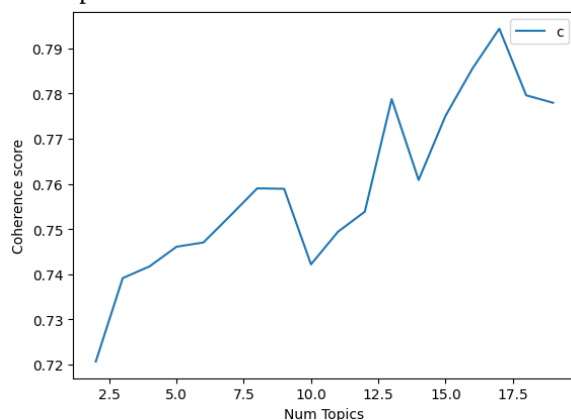
Sebelum	Sesudah
pju jalan raya satelit indah padam	['pju', 'jalan', 'raya', 'satelit', 'indah', 'padam']

4.4. Topic Modeling

Tahapan *topic modelling* dilakukan untuk menganalisis tema utama dari data keluhan warga Surabaya menggunakan metode *Latent Dirichlet Allocation* (LDA) dengan *library Gensim*. Proses dimulai dengan menggunakan data yang telah melalui tokenisasi, kemudian diubah menjadi format daftar.

Selanjutnya, dilakukan pembangunan model bigram dan trigram untuk mengidentifikasi frasa penting. Proses TF-IDF juga dilakukan untuk menyaring kata-kata yang relevan, yaitu kata-kata yang memiliki bobot signifikan. Hasil dari tahapan ini digunakan untuk membentuk dictionary dan corpus. Dictionary merupakan kumpulan kumpulan dari kata kata unik yang didapat dari keseluruhan data. Sedangkan corpus merupakan representasi dokumen dalam bentuk pasangan ID kata dari dictionary dan bobot atau frekuensinya dalam dokumen. Filtering dilakukan pada *dictionary* dan *corpus* untuk membuang kata yang terlalu umum atau jarang muncul, sehingga hanya menyisakan kata yang signifikan untuk analisis topik.

Conherence score kemudian dihitung berdasarkan jumlah topik yang diuji. Variasi parameter seperti {min_count}, {threshold}, {max_df}, {min_df}, {no_below}, {no_above}, dan {random_state} diuji untuk menemukan hasil optimal. Kombinasi terbaik ditemukan pada min_count=3, threshold=10, max_df=0.5, min_df=150, no_below=300, no_above=0.1, dan random_state=56, menghasilkan *coherence score* 0.794 dengan 17 topik. Grafik yang menggambarkan hasil perhitungan *coherence score* berdasarkan jumlah topik (*Num Topics*) yang diuji dalam pemodelan topik dapat dilihat pada Gambar 4.



Gambar 4. Gambar hasil perhitungan *conherence score*

Nilai *coherence score* meningkat seiring bertambahnya jumlah topik, dengan puncaknya berada di *Num Topic* = 17, menghasilkan *coherence score* sebesar 0.794, yang merupakan nilai terbaik dalam percobaan ini. Nilai tersebut menunjukkan bahwa model LDA menghasilkan koherensi antar topik yang sangat baik pada jumlah topik tersebut, sehingga menjadi pilihan yang optimal untuk analisis topik. Grafik hasil perhitungan *coherence score* berdasarkan jumlah topik dapat dilihat pada Gambar 4.

Setelah jumlah topik optimal ditentukan, dilakukan analisis lebih lanjut terhadap topik-topik yang dihasilkan oleh model. Setiap topik direpresentasikan oleh kumpulan kata-kata dengan bobot tertinggi yang mencerminkan tema utama dari topik tersebut. Kata-kata dominan ini digunakan sebagai dasar untuk memberikan interpretasi dan penamaan pada masing-masing topik, sehingga mempermudah identifikasi aspek-aspek utama dalam data keluhan.

Berdasarkan hasil analisis pada 17 topik, dilakukan penentuan nama topik berdasarkan kumpulan kata-kata yang ada pada masing-masing topik. Interpretasi aspek yang dihasilkan dari topik ke-1 hingga topik ke-17 dapat dilihat pada Tabel 9, yang memberikan gambaran tema atau isu utama yang muncul dalam data.

Tabel 9. Hasil interpretasi aspek

No.	Nama Aspek
1	Proses Administrasi dan Informasi Publik
2	Pengaduan dan Penyelesaian Keluhan
3	Permintaan Dan Pendaftaran Administrasi
4	Kualitas Pelayanan Masyarakat
5	Pendidikan dan Sekolah
6	Durasi dan Efisiensi Layanan
7	Keluhan Kebutuhan Lapangan Pekerjaan
8	Masalah Pohon dan Gangguan Lingkungan
9	Kerusakan Fasilitas dan Infrastruktur
10	Kondisi Fisik Lingkungan
11	Permintaan Informasi dan Bantuan
12	Pelaporan dan Tindak Lanjut Masalah
13	Permohonan Perbaikan Dan Pembaharuan
14	Masalah di Lingkungan Masyarakat
15	Lamanya Proses Pengajuan
16	Permasalahan Parkir
17	Pemimpin Wilayah yang Bermasalah

4.5. Klasifikasi Sentimen

Setelah pemodelan topik, tahap klasifikasi sentimen dilakukan menggunakan algoritma *Multinomial Naïve Bayes* (MNB). Sebelum melakukan pelatihan model, data teks diubah terlebih dahulu menjadi representasi numerik menggunakan TF-IDF, untuk merepresentasikan bobot kata berdasarkan frekuensi lokal dan global dalam dataset.

Selanjutnya, beberapa skenario pengujian dilakukan untuk mengidentifikasi kombinasi terbaik dari *preprocessing*, metode sampling, dan pembagian data untuk menemukan akurasi yang paling optimal.

Berikut adalah rangkuman langkah-langkah dan hasil pengujian:

4.5.1. Pengolahan Dataset

Dataset yang digunakan telah melalui preprocessing, termasuk variasi penghapusan stopword dan stemming, untuk mengamati pengaruh masing-masing terhadap akurasi model.

4.5.2. Pembagian Data

Pembagian data dilakukan dengan dua metode:

- Holdout*: Data dibagi dalam beberapa rasio, yaitu 70:30, 80:20, dan 90:10, untuk data latih dan uji.
- Cross-Validation*: Data dipecah menjadi 5-fold dan 10-fold untuk memastikan performa model yang lebih stabil.

4.5.3. Penanganan Ketidakseimbangan Kelas

Ketidakseimbangan kelas ditangani menggunakan dua teknik:

- SMOTE (Synthetic Minority Over-sampling Technique)*: Menambahkan sampel sintetis pada kelas minoritas.
- RUS (Random Under-Sampling)*: Mengurangi jumlah sampel pada kelas mayoritas.

4.5.4. Analisis Hasil Pengujian Skenario

Setelah model Naïve Bayes dilatih menggunakan kombinasi rasio data latih-uji, metode sampling, dan variasi preprocessing. Evaluasi dilakukan dengan metrik akurasi, precision, recall, dan F1-score untuk memahami performa model. Hasil pengujian skenario ini kemudian disajikan pada Tabel 10 untuk memberikan wawasan menyeluruh mengenai konfigurasi metode terbaik.

Tabel 10. Tabel analisis pengujian skenario

Stopword Removed			Akurasi (train)	Akurasi (test)	Overfit/ Underfit/ Good	Precision	Recall	F1-Score
Holdout	10:90	Normal	0.81	0.75	Good	0.78	0.70	0.71
		SMOTE	0.84	0.79	Good	0.78	0.78	0.78
		RUS	0.84	0.80	Good	0.79	0.79	0.79
	20:80	Normal	0.81	0.74	Good	0.77	0.70	0.70
		SMOTE	0.85	0.78	Good	0.77	0.78	0.77
		RUS	0.85	0.78	Good	0.77	0.78	0.78
	30:70	Normal	0.81	0.75	Good	0.77	0.70	0.71
		SMOTE	0.86	0.77	Good	0.76	0.77	0.77
		RUS	0.85	0.77	Good	0.77	0.77	0.77
Cross-Validation	5	Normal	0.75	0.73	Good	0.77	0.69	0.69
		SMOTE	0.85	0.75	Good	0.74	0.75	0.75
		RUS	0.85	0.75	Good	0.75	0.75	0.75
	10	Normal	0.76	0.73	Good	0.77	0.70	0.70
		SMOTE	0.85	0.75	Good	0.75	0.75	0.75
		RUS	0.85	0.75	Good	0.75	0.75	0.75
Stemmed			Akurasi (train)	Akurasi (test)	Overfit/ Underfit/ Good	Precision	Recall	F1-Score
Holdout	10:90	Normal	0.79	0.74	Good	0.76	0.70	0.70
		SMOTE	0.82	0.78	Good	0.78	0.77	0.78
		RUS	0.82	0.79	Good	0.78	0.78	0.78
	20:80	Normal	0.79	0.73	Good	0.76	0.69	0.69
		SMOTE	0.83	0.76	Good	0.76	0.76	0.76
		RUS	0.82	0.77	Good	0.76	0.76	0.76
	30:70	Normal	0.79	0.74	Good	0.77	0.70	0.71
		SMOTE	0.83	0.76	Good	0.75	0.76	0.75
		RUS	0.83	0.76	Good	0.75	0.76	0.76
Cross-Validation	5	Normal	0.73	0.72	Good	0.76	0.68	0.68
		SMOTE	0.83	0.74	Good	0.74	0.74	0.74
		RUS	0.83	0.75	Good	0.74	0.74	0.74
	10	Normal	0.75	0.73	Good	0.76	0.69	0.69
		SMOTE	0.83	0.74	Good	0.74	0.74	0.74
		RUS	0.82	0.75	Good	0.74	0.74	0.74

4.5.5. Hasil dan Pemilihan Skenario Terbaik

Berdasarkan hasil evaluasi, skenario terbaik menggunakan kombinasi preprocessing stopword removal, teknik penyeimbangan data RUS, dan rasio data 10:90 pada metode *Holdout*. Kombinasi ini

menghasilkan akurasi 0.84 pada data latih dan 0.80 pada data uji, serta *precision*, *recall*, dan *F1-score* masing-masing sebesar 0.79. Model menunjukkan performa optimal tanpa tanda-tanda *overfitting* atau *underfitting*. Hasil akhir dari prediksi klasifikasi Naïve

Bayes dengan skenario terbaik diintegrasikan ke dalam dataframe analisis sentimen untuk digunakan pada analisis berbasis aspek.

4.6. Analisis Sentimen Berbasis Aspek

Setelah topik-topik keluhan masyarakat telah berhasil diidentifikasi menggunakan pendekatan *Latent Dirichlet Allocation* (LDA), dan sentimen

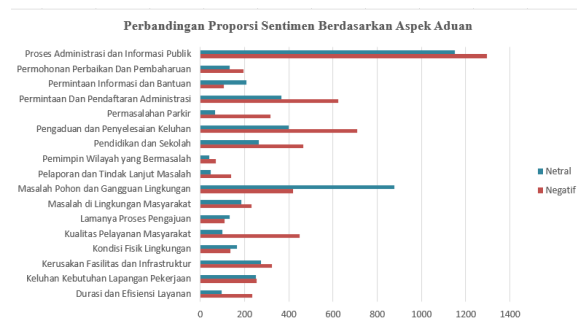
terhadap setiap topik telah diklasifikasikan menggunakan algoritma *Multinomial Naive Bayes* (MNB). Analisis sentiment berbasis aspek dilakukan. Distribusi jumlah sentimen negatif dan netral dalam berbagai aspek keluhan yang telah diidentifikasi dapat dilihat secara lebih rinci pada Tabel 11.

Tabel 11. Distribusi sentimen berdasarkan aspek aduan

Aspek	Negatif	Netral	Total
Durasi dan Efisiensi Layanan	234	95	329
Keluhan Kebutuhan Lapangan Pekerjaan	256	252	508
Kerusakan Fasilitas dan Infrastruktur	324	274	598
Kondisi Fisik Lingkungan	137	164	301
Kualitas Pelayanan Masyarakat	450	100	550
Masalah di Lingkungan Masyarakat	231	186	417
Masalah Pohon dan Gangguan Lingkungan	419	877	1296
Pelaporan dan Tindak Lanjut Masalah	140	47	187
Pendidikan dan Sekolah	464	266	730
Pengaduan dan Penyelesaian Keluhan	708	401	1109
Pemimpin Wilayah yang Bermasalah	71	42	113
Permasalahan Parkir	316	66	382
Permintaan Dan Pendaftaran	624	368	992
Permintaan Informasi dan Bantuan	107	210	317
Permohonan Perbaikan Dan Pembaharuan	195	132	327
Proses Administrasi dan Informasi Publik	1296	1152	2448
Lamanya Proses Pengajuan	110	134	244
Grand Total	6082	4766	10848

Berdasarkan Tabel 11, aspek dengan keluhan tertinggi adalah "Proses Administrasi dan Informasi Publik" dengan 2.448 keluhan, terdiri dari 1.296 keluhan negatif dan 1.152 netral, menunjukkan ketidakpuasan masyarakat terhadap isu administratif.

Aspek dengan keluhan terendah adalah "Penyelesaian Masalah dan Solusi Keluhan" dengan 113 keluhan, 71 di antaranya negatif. Aspek "Masalah Pohon dan Gangguan Lingkungan" memiliki dominasi sentimen netral tertinggi, sedangkan "Proses Administrasi dan Informasi Publik" dan "Pengaduan dan Penyelesaian Keluhan" didominasi oleh sentimen negatif, yang menunjukkan kebutuhan untuk perhatian lebih.



Gambar 5. Perbandingan proporsi sentimen berdasarkan aspek aduan

Pada Gambar 5 dapat dilihat grafik perbandingan proporsi sentimen negatif dan netral berdasarkan aspek keluhan.

5. KESIMPULAN DAN SARAN

Berdasarkan analisis sentimen berbasis aspek pada sistem layanan pengaduan masyarakat di Surabaya menggunakan metode *Latent Dirichlet Allocation* (LDA) dan *Naïve Bayes*, penelitian ini menemukan bahwa LDA berhasil mengidentifikasi 17 topik utama yang menggambarkan isu-isu keluhan masyarakat, seperti proses administrasi, pengaduan, kualitas pelayanan, dan masalah infrastruktur. Klasifikasi sentimen menggunakan *Naïve Bayes* mengkategorikan mayoritas sentimen sebagai negatif, khususnya terkait topik administrasi publik, yang menunjukkan ketidakpuasan masyarakat dan menjadi perhatian utama untuk perbaikan. Model *Naïve Bayes* menunjukkan performa optimal dengan akurasi 80%, dengan peningkatan yang dicapai melalui metode *resampling*. Penelitian ini menyarankan penggunaan data yang lebih besar dan beragam, serta penerapan metode klasifikasi yang lebih kompleks seperti SVM atau Random Forest untuk penelitian selanjutnya.

DAFTAR PUSTAKA

- [1] G. Thabrani, "Pelayanan Publik: Pengertian, Jenis, Prinsip, Dimensi, Indikator, dsb," serupa.id.
- [2] A. Kurniawan, "WargaKu, Ini Dia Aplikasi untuk Komunikasi Warga dengan Pemko Surabaya," SINDOnews Daerah.
- [3] Dinas Komunikasi Dan Informatika Surabaya, "Analisa Kepuasan Masyarakat Terhadap

- Program Pelayanan Pengaduan Media Center,” 2019.
- [4] S. Roiqoh and B. Zaman, “JURNAL MEDIA INFORMATIKA BUDIDARMA Analisis Sentimen Berbasis Aspek Ulasan Aplikasi Mobile JKN dengan Lexicon Based dan Naïve Bayes,” 2023, doi: 10.30865/mib.v7i3.6194.
- [5] C. H. Yutika, A. Adiwijaya, and S. Al Faraby, “Analisis Sentimen Berbasis Aspek pada Review Female Daily Menggunakan TF-IDF dan Naïve Bayes,” *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 5, no. 2, p. 422, Apr. 2021, doi: 10.30865/mib.v5i2.2845.
- [6] K. T. Pamungkas, L. Aridinanti, and W. Wibowo, “Analisis Sentimen Pelaporan Masyarakat di Situs Media Centre Surabaya dengan Naïve Bayes Classifier,” *JURNAL TEKNIK ITS*, vol. 11, no. 2, pp. 197–203, 2022.
- [7] Dinas Komunikasi Dan Informatika Surabaya, “Indeks Kepuasan Masyarakat Aplikasi Wargaku Surabaya,” 2022.
- [8] F. A. D. P. Febrianti, F. Hamami, and R. Y. Fa’rifah, “ASPECT-BASED SENTIMENT ANALYSIS TERHADAP ULASAN APLIKASI FLIP MENGGUNAKAN PEMBOBOTAN TERM FREQUENCY-INVERSE DOCUMENT FREQUENCY (TF-IDF) DENGAN METODE KLASIFIKASI K-NEAREST NEIGHBORS (K-NN),” *Jurnal Indonesia : Manajemen Informatika dan Komunikasi*, vol. 4, no. 3, pp. 1858–1873, Sep. 2023, doi: 10.35870/jimik.v4i3.429.
- [9] H. Liu, I. Chatterjee, M. Zhou, X. S. Lu, and A. Abusorrah, “Aspect-Based Sentiment Analysis: A Survey of Deep Learning Methods,” Dec. 01, 2020, *Institute of Electrical and Electronics Engineers Inc.* doi: 10.1109/TCSS.2020.3033302.
- [10] D. A. Kristiyanti and Sri Hardani, “Sentiment Analysis of Public Acceptance of Covid-19 Vaccines Types in Indonesia using Naïve Bayes, Support Vector Machine, and Long Short-Term Memory (LSTM),” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 7, no. 3, pp. 722–732, Jun. 2023, doi: 10.29207/resti.v7i3.4737.
- [11] A. Sabrani, I. W. Gede Putu Wirarama Wedashwara, and F. Bimantoro, “METODE MULTINOMIAL NAÏVE BAYES UNTUK KLASIFIKASI ARTIKEL ONLINE TENTANG GEMPA DI INDONESIA (Multinomial Naïve Bayes Method for Classification of Online Article About Earthquake in Indonesia).” [Online]. Available: <http://jtika.if.unram.ac.id/index.php/JTIKA/>
- [12] W. A. N. Sari and H. D. Purnomo, “TOPIC MODELING USING THE LATENT DIRICHLET ALLOCATION METHOD ON WIKIPEDIA PANDEMIC COVID-19 DATA IN INDONESIA,” *Jurnal Teknik Informatika (Jutif)*, vol. 3, no. 5, pp. 1223–1230, Oct. 2022, doi: 10.20884/1.jutif.2022.3.5.321.
- [13] J. C. Campbell, A. Hindle, and E. Stroulia, “Latent Dirichlet Allocation: Extracting Topics from Software Engineering Data,” 2014.