

ANALISIS TOPIK DAN SENTIMEN PADA ULASAN PASIEN MENGGUNAKAN K-MEANS CLUSTERING DAN LEXICON-BASED SCORING (STUDI KASUS: RUMAH SAKIT XYZ)

Irsyad Nurhakim^{1*}, Ade Andri Hendriadi², Ahmad Khusaeri³

^{1,2,3}Universitas Singaperbangsa Karawang; Jl. HS. Ronggo Waluyo, Puserjaya, Telukjambe Timur, Karawang, Jawa Barat, 41361; (0267) 641177

Keywords:

Text Mining;
Analisis Sentimen;
K-Means Clustering;
Lexicon-Based;
Ulasan Pasien.

Correspondent Email:

irsyadnurhakim816@gmail.com

Abstrak. Penelitian ini bertujuan memetakan persepsi kualitas layanan Rumah Sakit XYZ melalui analisis 2.324 ulasan pasien di Google Maps. Teks tidak terstruktur diolah menggunakan pendekatan *text mining* yang menggabungkan *Lexicon-Based Scoring* (kamus InSet) untuk analisis sentimen dan algoritma *K-Means Clustering* berbasis TF-IDF untuk pemodelan topik. Hasil analisis menunjukkan mayoritas ulasan bersentimen positif (83,35%) dan sisanya negatif (16,65%). Algoritma K-Means menghasilkan empat kluster topik optimal: Administrasi & BPJS (60,33%), Pelayanan Tenaga Medis (24,57%), Layanan Rawat Inap (10,54%), serta Fasilitas & Kenyamanan (4,56%). Integrasi kedua metode mengungkap bahwa Administrasi & BPJS adalah area paling kritis dengan ketidakpuasan tertinggi (26,2%). Sebagai implikasi manajerial, penelitian merekomendasikan integrasi SIMRS dengan Mobile JKN dan *E-Prescription* guna mereduksi waktu antrean dan mendukung peningkatan mutu layanan berbasis data.



Copyright © 2026 The Author(s),
Published by Jurnal Informatika dan
Teknik Elektro Terapan. This is an
open-access article distributed under
the terms of the Creative Commons
Attribution-NonCommercial 4.0
International License (CC BY-NC
4.0).

Abstract. This study maps the perception of service quality at XYZ Hospital by analyzing 2,324 patient reviews on Google Maps. Unstructured text was processed using a text mining approach combining *Lexicon-Based Scoring* (InSet dictionary) for sentiment analysis and TF-IDF-based *K-Means Clustering* for topic modeling. The results showed the majority of reviews had positive sentiment (83.35%), while the rest were negative (16.65%). The K-Means algorithm generated four optimal topic clusters: Administration & BPJS (60.33%), Medical Staff Services (24.57%), Inpatient Services (10.54%), and Facilities & Comfort (4.56%). Integrating both methods revealed that Administration & BPJS is the most critical area with the highest dissatisfaction rate (26.2%). For managerial implications, the study recommends integrating SIMRS with Mobile JKN and *E-Prescription* to reduce waiting times and support data-driven service quality improvement.

1. PENDAHULUAN

Kemajuan teknologi digital dalam satu dekade terakhir telah membawa perubahan signifikan terhadap pola perilaku masyarakat, khususnya dalam mengevaluasi kualitas layanan publik. Saat ini, masyarakat

mengandalkan internet sebagai sumber utama memperoleh informasi sebelum menentukan pilihan layanan kesehatan [1]. Studi di Indonesia mencatat bahwa 98,4% responden menggunakan internet untuk mencari informasi kesehatan, yang dianggap krusial dalam

pengambilan keputusan [2]. Dinamika ini memunculkan *User-Generated Content* (UGC), yakni konten ulasan langsung dari pengguna pada platform digital [3]. Google Maps berperan signifikan sebagai media masyarakat untuk mengevaluasi layanan yang diterima [4]. Suara pelanggan kini semakin transparan dan berpengaruh substansial dalam membentuk citra institusi pelayanan kesehatan.

Dalam sektor kesehatan, kepercayaan adalah fondasi utama. Calon pasien menelaah ulasan daring sebelum memilih rumah sakit [1]. Literatur menunjukkan bahwa 37% pasien akan membatalkan niat memilih penyedia layanan jika menemukan ulasan negatif [5]. Lebih lanjut, peringkat kualitas rumah sakit terbukti memiliki pengaruh lebih dominan terhadap keputusan pasien dibandingkan peringkat dokter secara individu [6]. Ulasan daring memiliki keterkaitan signifikan dengan persepsi masyarakat terhadap dimensi medis maupun nonmedis [7]. Reputasi digital kini berfungsi sebagai aset strategis bagi manajemen rumah sakit untuk mempertahankan citra positif.

Namun, peningkatan volume ulasan daring secara masif menghasilkan *big data* berupa teks tidak terstruktur. Proses analisis manual menjadi tidak efisien, memakan waktu, dan memicu bias subjektif [8]. Tanpa pendekatan komputasional, institusi kesulitan mengidentifikasi pola keluhan, apresiasi, maupun saran layanan [4]. Kondisi ini menimbulkan kesenjangan antara data yang melimpah dan kemampuan mengolahnya menjadi informasi kuantitatif yang terstruktur, yang pada akhirnya berpotensi menghambat pengambilan keputusan berorientasi data (*data-driven decision making*) [7], [8]. Padahal, pemanfaatan sistem analisis data yang terintegrasi sangat esensial untuk meningkatkan efisiensi operasional, mempercepat pengambilan keputusan manajerial, serta mengoptimalkan mutu layanan kesehatan bagi pasien [9].

Text Mining hadir sebagai solusi analitis untuk mengekstraksi pola, topik, dan sentimen dari informasi tidak terstruktur [10]. Penelitian ini menerapkan kombinasi metode *K-Means Clustering* dan *Lexicon-Based Scoring*. *K-Means* dipilih karena keunggulan efisiensi komputasi untuk mengelompokkan dokumen teks tanpa label [11], yang banyak diimplementasikan di Indonesia [8], [12], dan

terbukti efektif menghasilkan identifikasi topik akurat [13]. Di sisi lain, *Lexicon-Based Scoring* digunakan untuk mengukur sentimen dengan memberikan bobot nilai kata [10]. Pendekatan ini, khususnya menggunakan kamus InSet, andal menangani teks bahasa Indonesia tanpa memerlukan proses pelatihan model (*training data*) yang rumit [14]. Sinergi kedua metode ini memfasilitasi analisis dua dimensi: identifikasi topik dan kecenderungan perasaan pasien.

Pendekatan ini diterapkan pada studi kasus Rumah Sakit XYZ, rumah sakit swasta rujukan utama di wilayah perkotaan. Rumah sakit ini berada dalam fase peningkatan standar klasifikasi yang berdampak signifikan pada mutu layanan [15]. Transisi ini memicu dinamika ekspektasi publik seiring kecenderungan masyarakat yang mengandalkan ulasan daring [2], [6]. Oleh karena itu, penelitian ini bertujuan untuk mengelompokkan data teks ulasan pasien dari Google Maps menggunakan *K-Means Clustering* untuk menemukan pola topik utama, serta menganalisis tingkat kepuasan menggunakan metode *Lexicon-Based Scoring*. Hasil integrasi analisis tersebut diarahkan untuk menyusun pemetaan prioritas perbaikan area layanan, yang selanjutnya direkomendasikan kepada pihak manajemen sebagai dukungan keputusan strategis demi peningkatan kualitas layanan kesehatan secara berkelanjutan.

2. TINJAUAN PUSTAKA

2.1. Text Mining

Text mining merupakan pendekatan komputasional dalam *Knowledge Discovery in Database* (KDD) yang berfokus pada ekstraksi pengetahuan tersembunyi dari teks tidak terstruktur melalui tahapan prapemrosesan, representasi numerik, dan analisis pemodelan [16]. Dalam konteks evaluasi layanan kesehatan, sumber data *text mining* yang sangat strategis adalah *User-Generated Content* (UGC) seperti ulasan Google Maps. Ulasan daring ini merefleksikan pengalaman empiris serta persepsi independen pasien terhadap mutu layanan klinis maupun non-klinis dari sebuah rumah sakit [1], [17].

2.2. TF-IDF (Term Frequency-Inverse Document Frequency)

Sebelum dianalisis oleh algoritma, teks hasil prapemrosesan harus ditransformasikan ke dalam representasi vektor numerik (*feature extraction*). Penelitian ini menggunakan TF-IDF yang mengintegrasikan frekuensi kemunculan kata dalam satu dokumen ($tf_{i,j}$) dan kemunculan kata tersebut pada keseluruhan korpus (df_i). Pembobotan ini memastikan bahwa kata yang spesifik dan bermakna memiliki nilai yang lebih tinggi dibandingkan kata umum [18], [19]. Secara sistematis, bobot TF-IDF ($W_{i,j}$) dirumuskan sebagai:

$$W_{i,j} = tf_{i,j} \times \log\left(\frac{N}{df_i}\right) \quad (1)$$

Di mana N adalah jumlah total dokumen ulasan.

2.3. K-Means Clustering, Elbow, dan Silhouette

K-Means merupakan algoritma *unsupervised learning* yang mempartisi dokumen ke dalam sejumlah k kluster berdasarkan nilai kedekatan jarak terhadap pusat kluster (*centroid*) [12]. Fungsi objektif algoritma ini dirumuskan sebagai [20]:

$$E = \sum_{i=1}^k \sum_{p \in C_i} |p - m_i|^2 \quad (2)$$

Penentuan jumlah kluster (k) yang paling optimal dievaluasi menggunakan Elbow Method berdasarkan titik penurunan ekstrem pada nilai *Within-Cluster Sum of Squares* (WCSS). Kualitas hasil pengelompokan divalidasi secara kuantitatif melalui *Silhouette Score* (s) yang dirumuskan sebagai:

$$s = \frac{b - a}{\max(a, b)} \quad (3)$$

Di mana a merepresentasikan rata-rata jarak intra-kluster dan b adalah rata-rata jarak inter-kluster terdekat. Nilai s berada pada rentang -1 hingga +1, di mana nilai yang mendekati +1 menandakan pengelompokan yang optimal [21]. Hasil ini kemudian divisualisasikan melalui metode reduksi dimensi *Principal Component Analysis* (PCA).

2.4. Lexicon-Based Scoring

Lexicon-Based Scoring adalah metode analisis sentimen yang menentukan polaritas teks bersandar pada nilai bawaan leksikal dari

sebuah kamus, tanpa memerlukan data latih [22]. Penelitian ini menggunakan *Indonesian Sentiment Lexicon* (InSet) [23]. Skor sentimen total suatu dokumen dihitung dengan menjumlahkan seluruh bobot kata yang cocok dengan entri kamus:

$$Total\ Score = \sum_{i=1}^n v_i \quad (4)$$

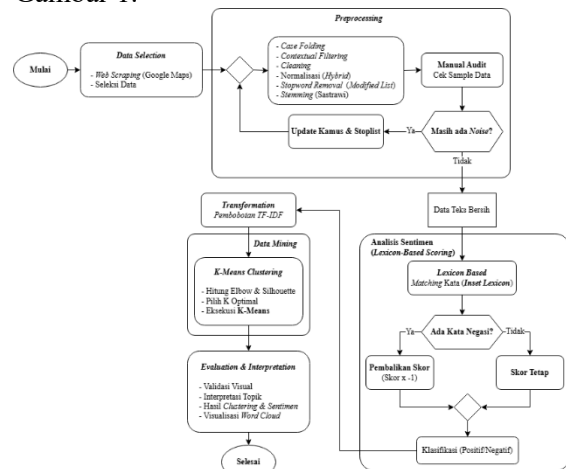
Di mana v_i adalah nilai bobot sentimen dari kata ke- i , dan n adalah jumlah kata yang cocok. Berdasarkan total skor tersebut, dokumen diklasifikasikan menjadi tiga kelas: Positif ($Total\ Score > 0$), Negatif ($Total\ Score < 0$), dan Netral ($Total\ Score = 0$) [10]. Guna menjaga akurasi konteks, algoritma ini dimodifikasi dengan penanganan negasi (*negation handling*). Skor akhir suatu kata (S) dihitung melalui persamaan:

$$S = S(N) \times S(Lex) \quad (5)$$

Di mana $S(Lex)$ adalah bobot sentimen dasar dari kamus, dan $S(N)$ bernilai -1 jika kata sifat didahului oleh penyangkal (seperti "tidak", "bukan", atau "kurang"), serta bernilai 1 jika tidak ada negasi [24].

3. METODE PENELITIAN

Rancangan penelitian ini disusun berdasarkan tahapan *Knowledge Discovery in Databases* (KDD) yang disesuaikan dengan kebutuhan analisis text mining dan pengolahan ulasan rumah sakit. Alur rancangan penelitian secara keseluruhan divisualisasikan pada Gambar 1.



Gambar 1 Alur Rancangan Penelitian

3.1. Data Selection

Tahap seleksi data diawali dengan akuisisi ulasan Rumah Sakit XYZ dari platform Google Maps menggunakan teknik *web scraping* berbasis Selenium dan Python. Penelitian ini menerapkan metode *Purposive Sampling* untuk memastikan relevansi data, di mana kriteria inklusi difokuskan pada ulasan yang memuat narasi teks (*textual review*) dan menggunakan Bahasa Indonesia. Sistem secara otomatis mengeliminasi ulasan yang hanya berisi peringkat bintang (*rating*) tanpa deskripsi. Tahap ini menghasilkan himpunan data mentah (*raw data*) yang siap diproses pada tahap prapemrosesan.

3.2. Preprocessing

Data teks mentah hasil seleksi diproses melalui serangkaian tahapan prapemrosesan untuk menyiapkan data sebelum analisis. Penelitian ini menerapkan teknik *preprocessing* yang disesuaikan secara kontekstual dengan karakteristik ulasan layanan rumah sakit, guna meminimalkan potensi bias pada hasil *clustering* serta menjaga konsistensi pengukuran sentimen. Tahapan yang dilakukan meliputi:

1. *Case Folding*: Mengubah seluruh karakter huruf dalam dokumen teks menjadi huruf kecil (*lowercase*) untuk menyeragamkan format penulisan.
2. *Contextual Filtering*: Melakukan penghapusan atau normalisasi pada frasa spesifik (seperti nama instansi rumah sakit) yang berpotensi menjadi *domain-specific stopwords* dan menimbulkan bias dominasi *centroid* pada proses *clustering*.
3. *Cleaning*: Menghapus elemen gangguan (*noise*) yang tidak diperlukan dalam analisis, seperti tanda baca, angka, simbol, emotikon, serta tautan (URL).
4. Normalisasi: Mentransformasi kata tidak baku, singkatan, akronim, dan bahasa informal menjadi bentuk baku. Penelitian ini menggunakan pendekatan hibrida dengan mengintegrasikan Kamus Eksternal (*Colloquial Indonesian Lexicon* [25]) dan Kamus Domain Spesifik (seperti "ranap" menjadi "rawat inap") yang disusun melalui audit frekuensi kata.
5. *Stopword Removal*: Menghapus kata-kata umum yang tidak memiliki bobot informasi pembeda menggunakan pustaka

Sastrawi yang dimodifikasi. Modifikasi kritis dilakukan melalui pengecualian kata negasi (seperti "tidak", "bukan", "kurang") untuk menjaga konteks semantik sentimen, serta penambahan daftar hapus untuk kata berfrekuensi sangat tinggi (seperti "layan", "sangat").

6. *Stemming*: Mengembalikan kata berimbuhan ke bentuk kata dasar (*root word*) menggunakan algoritma Nazief-Adriani melalui pustaka Sastrawi untuk menyatukan variasi kata yang memiliki makna semantik yang sama.
7. Audit Manual: Evaluasi iteratif pasca-*stemming* terhadap sampel data untuk mengidentifikasi kata tidak baku yang belum tertangani. Jika ditemukan pola baru, kamus diperbarui dan proses prapemrosesan dijalankan kembali hingga data dinilai bersih dan konsisten.

3.3. Analisis Sentimen (Lexicon-Based Scoring)

Tahap analisis sentimen bertujuan untuk mengekstraksi polaritas opini serta mengukur tingkat kepuasan pasien. Analisis dilakukan menggunakan pendekatan *lexicon-based* dengan mengacu pada kamus InSet (*Indonesian Sentiment Lexicon*) [23]. Penelitian ini menerapkan teknik Penanganan Negasi (*Negation Handling*) guna menangani potensi kesalahan interpretasi polaritas. Sistem mendeteksi kata negasi yang muncul sebelum kata sifat; apabila kondisi tersebut terpenuhi, maka skor polaritas kata sifat akan dibalik dengan cara mengalikan nilai polaritas dengan -1.

Skor sentimen total pada setiap ulasan diperoleh melalui akumulasi seluruh skor kata setelah penyesuaian negasi. Selanjutnya, ulasan diklasifikasikan ke dalam kategori sentimen positif ($\text{skor} > 0$) dan negatif ($\text{skor} < 0$). Ulasan dengan skor 0 (netral) dieliminasi dan tidak dilanjutkan pada tahap analisis berikutnya. Mempertahankan data netral berisiko mendilusi pembobotan fitur pada matriks TF-IDF, sehingga kluster yang dihasilkan akan dipenuhi noise informasi umum. Keandalan hasil klasifikasi yang dilakukan secara otomatis oleh sistem ini kemudian divalidasi secara kualitatif oleh pakar bahasa Indonesia.

3.4. Transformation (TF-IDF)

Ulasan yang memiliki polaritas positif dan negatif ditransformasikan ke dalam bentuk representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Teknik ini menghitung bobot setiap kata berdasarkan frekuensi kemunculan kata dalam suatu dokumen dan tingkat keunikan kata tersebut dalam keseluruhan korpus. Pendekatan ini memungkinkan kata yang bermakna spesifik memperoleh bobot yang lebih tinggi. Proses transformasi menghasilkan matriks TF-IDF berdimensi tinggi yang menjadi input numerik pada algoritma pemodelan topik.

3.5. Data Mining (K-Means Clustering)

Tahap ekstraksi pengetahuan dilakukan melalui pengelompokan topik menggunakan algoritma *K-Means Clustering*. Algoritma ini mempartisi representasi numerik TF-IDF ke dalam sejumlah kluster berdasarkan kedekatan jarak (*Euclidean distance*) terhadap titik pusat (*centroid*). Penentuan jumlah kluster optimal (*k*) dievaluasi melalui dua pendekatan:

1. *Elbow Method*: Pengujian variasi jumlah kluster dengan menghitung nilai inersia (*Within-Cluster Sum of Squares / WCSS*) untuk mengidentifikasi "titik siku" (*elbow point*) penurunan variasi.
2. Validasi *Silhouette Score*: Mengukur tingkat kohesi intra-kluster dan separasi antar-kluster untuk memastikan keseimbangan optimal antara kompleksitas model dan homogenitas kluster.

3.6. Evaluation & Interpretation

Tahap akhir bertujuan mengevaluasi struktur kluster dan menginterpretasikan hasil penelitian. Evaluasi struktur kluster yang berdimensi tinggi dilakukan secara visual menggunakan teknik reduksi dimensi *Principal Component Analysis* (PCA). Selanjutnya, interpretasi makna setiap kluster dilakukan dengan menganalisis nilai *centroid* untuk mengidentifikasi kata-kata dengan bobot TF-IDF tertinggi sebagai fitur yang paling representatif untuk pelabelan nama topik.

Hasil pengelompokan topik (K-Means) kemudian diintegrasikan dengan polaritas sentimen (*Lexicon-Based*) pada tingkat *record* dokumen. Teknik tabulasi silang (*cross-tabulation*) diterapkan untuk menghitung

proporsi distribusi sentimen positif dan negatif di dalam masing-masing topik layanan. Terakhir, visualisasi *Word Cloud* digunakan untuk memberikan representasi grafis terhadap frekuensi dan dominasi kata kunci utama yang menyusun persepsi ulasan pasien.

4. HASIL DAN PEMBAHASAN

4.1. Data Selection

Tahap awal dalam siklus KDD adalah data selection, yang bertujuan untuk menentukan himpunan data ulasan pasien yang relevan. Proses ini diawali dengan akuisisi data dari laman Google Maps Rumah Sakit XYZ menggunakan teknik *web scraping* berbasis Selenium.

Proses *scraping* dieksekusi pada 30 November 2025 dengan menerapkan teknik *Purposive Sampling*. Sistem dikonfigurasi untuk memfilter ulasan secara otomatis dengan kriteria inklusi: ulasan wajib memiliki narasi teks berbahasa Indonesia dan mengeliminasi ulasan yang hanya memberikan peringkat bintang (*rating*).

Dari proses tersebut diperoleh sebanyak 2.460 data mentah (*raw data*) dengan rekam jejak digital sejak tahun 2013 hingga 30 November 2025. Tren peningkatan volume ulasan dari tahun ke tahun disajikan pada Gambar 2.



Gambar 2 Tren Ulasan Mentah per Tahun

Berdasarkan Gambar 2, terlihat adanya tren peningkatan volume ulasan yang sangat signifikan. Pada periode awal (2013-2018), partisipasi pasien masih sangat minim, namun terjadi lonjakan tajam mulai tahun 2022 hingga mencapai puncaknya pada 2025 (956 ulasan). Dominasi data pada rentang waktu 2022-2025 ini memastikan bahwa korpus teks secara alamiah lebih merepresentasikan kondisi dan kualitas layanan mutakhir rumah sakit.

Struktur data mencakup atribut identitas pengguna, peringkat, waktu, dan isi ulasan. Untuk menjaga etika penelitian, identitas pengguna disamarkan melalui proses anonymization. Sampel data mentah hasil akuisisi ditunjukkan pada Tabel 1.

Tabel 1 Sampel Data Mentah Ulasan Pasien

ID Responden	Rating	Waktu	Ulasan (Review)
Resp_1	5	09/11/25	Tgl 6-7Oct, suami pertama kali dirawat di ruang ICU pasca pasang ring; bpjs; pelayanan dan perawatan tim medis di ICU Lt 3 yang fast response dan ramah; tidak jutek seperti RS lain sekitarnya yg pernah saya kunjungi. Kebersihan ruang icu ...
Resp_2	5	23/11/25	Saya pasien thalasemia. Pelayanan thalasemia di RS XYZ paling nyaman karena tanpa harus antri daftar😊Perawatnya baik, ramah, dokternya baik, adminnya baik, Penanganannya cepat, ruangnya sangat nyaman. Trimakasih sudah mempermudah kami untuk berobat RS XYZ 😊❤️

Berdasarkan Tabel 1, terlihat bahwa data masih bersifat tidak terstruktur (*unstructured text*) dan mengandung banyak *noise* seperti emotikon serta singkatan tidak baku. Oleh karena itu, kolom ulasan ditetapkan sebagai variabel input utama untuk diproses pada tahap prapemrosesan.

4.2. Preprocessing

Setelah tahap seleksi data, 2.460 ulasan mentah diproses karena masih mengandung elemen gangguan (*noise*) seperti inkonsistensi kapitalisasi, simbol, dan singkatan tidak baku. Penelitian ini menerapkan pipeline prapemrosesan khusus yang disesuaikan dengan karakteristik teks layanan rumah sakit. Proses ini bersifat sekuensial sekaligus iteratif, yang terdiri dari tujuh tahapan utama:

1. *Case Folding*: Menyeragamkan seluruh karakter menjadi huruf kecil (*lowercase*) untuk mencegah duplikasi fitur akibat perbedaan kapitalisasi.

2. *Contextual Filtering*: Mengeliminasi entitas spesifik seperti penyebutan nama rumah sakit ("RS XYZ") yang berstatus sebagai domain-specific stopwords. Jika tidak dihapus, kata ini akan mendominasi centroid kluster dan mengaburkan ekstraksi isu asli pasien.
3. *Cleaning*: Menggunakan pendekatan *Regular Expression* (Regex) untuk menghapus seluruh karakter non-alfabet, termasuk angka, tanda baca, dan emotikon.
4. *Normalisasi*: Menyeragamkan bahasa informal (*slang*) dan istilah medis menggunakan pendekatan hibrida. Contohnya, mengonversi "icu" menjadi "unit perawatan intensif" dan "lgsg" menjadi "langsung".
5. *Stopword Removal*: Menghapus kata umum tak bermakna menggunakan pustaka Sastrawi yang dimodifikasi. Modifikasi kritis dilakukan dengan mempertahankan kata negasi (seperti "tidak", "bukan", "kurang") agar konteks sentimen tidak hilang, serta menambahkan stopwords khusus untuk kata berfrekuensi sangat tinggi (seperti "pelayanan" dan "sangat").
6. *Stemming*: Mengembalikan kata berimbuhan ke bentuk dasar (misalnya "dirawat" menjadi "rawat") untuk menyederhanakan ruang fitur (*feature space*) dan mengurangi redundansi pada komputasi TF-IDF.
7. *Audit Manual*: Evaluasi iteratif terhadap sampel data untuk mengidentifikasi token *noise* yang terlewat. Contohnya, menambahkan kata "suster" untuk diseragamkan menjadi "perawat" pada kamus normalisasi. Iterasi dilakukan hingga distribusi token stabil.

Setelah pipeline ini selesai dieksekusi dan dievaluasi, sejumlah ulasan yang menjadi kosong akibat pembersihan teks dieliminasi, sehingga menghasilkan 2.417 dokumen ulasan bersih. Contoh transformasi teks dari data mentah hingga menjadi bentuk *root word* yang siap dianalisis disajikan pada Tabel 2.

Tabel 2 Hasil Transformasi Teks Pada Tahap Prapemrosesan

Data Mentah (Awal)	Hasil Akhir Prapemrosesan (Bersih)
Tgl 6-7Oct, suami pertama kali dirawat di ruang ICU pasca pasang ring; bpjs; pelayanan dan perawatan tim medis di ICU Lt 3 yang fast response dan ramah; tidak jutek seperti RS lain sekitarnya yg pernah saya kunjungi. Kebersihan ruang icu ...	tanggal oktober suami pertama kali rawat ruang unit awat intensif pasca pasang ring bpjs awat tim medis unit awat intensif lantai cepat respons ramah tidak kasar sekitar pernah kunjung bersih ruang unit awat intensif
Saya pasien thalasemia. Pelayanan thalasemia di RS XYZ paling nyaman karena tanpa harus antri daftar ☺ Perawatnya baik, ramah, dokternya baik, adminnya baik, Penanganannya cepat, ruangnya sangat nyaman. Trimakasih sudah mempermudah kami untuk berobat RS XYZ ☺❤	pasien thalasemia thalasemia nyaman antri daftar awat ramah administrasi tangan cepat ruang nyaman mudah obat

4.3. Analisis Sentimen (Lexicon-Based Scoring)

Tahap analisis sentimen dilakukan terhadap 2.417 dokumen bersih menggunakan pendekatan *Lexicon-Based Scoring* dengan kamus InSet. Sistem menghitung akumulasi bobot kata dan secara khusus mengimplementasikan mekanisme negation handling. Kata negasi seperti "tidak", "bukan", dan "kurang" tidak dihitung sebagai skor mandiri, melainkan berfungsi membalikkan polaritas (dikalikan -1) dari kata sentimen yang mengikutinya.

Sebagai pembuktian komputasi, contoh penerapan *negation handling* dan hasil akumulasi skor akhir pada dokumen ulasan pasien disajikan pada Tabel 3.

Tabel 3. Perhitungan Skor Sentimen Dan Penanganan Negasi

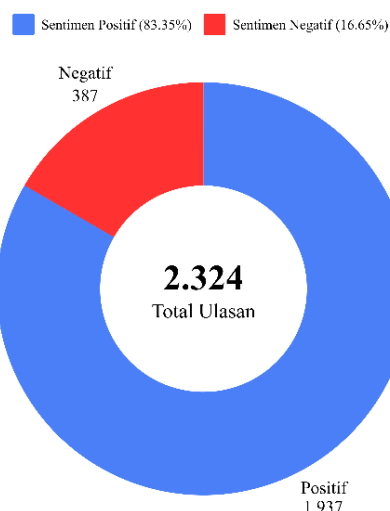
Ulasan Bersih	Rincian Bobot Kata dan Negasi	Skor	Label
rujuk tindak operasi nginap malam asa layan tugas awat ramah selalu beri solusi tenang alhamdulillah puas	operasi(+1), malam(-1), asa(+3), layan(+2), tugas(-1), awat(+5), ramah(+5), beri(+2), solusi(+3), tenang(+5), alhamdulillah(+5), puas(+3)	+32	Positif
keluarga pasien sakit ruang unit awat intensif ruang tunggu desain kurang	keluarga(+4), sakit(-5), ruang(-3), unit(-4), awat(+5), intensif(+2), desain(+4),	-8	Negatif

Ulasan Bersih	Rincian Bobot Kata dan Negasi	Skor	Label
sensitif tunggu butuh istirahat tidak sakit harus ruang tunggu	kurang(NEG) × sensitif(+1) = -1, butuh(-5), istirahat(-2), tidak(NEG) × sakit(-5) = +5, harus(-5)		

Berdasarkan akumulasi skor akhir tersebut, sistem melakukan klasifikasi dengan aturan: ulasan berskor > 0 dilabeli positif, < 0 dilabeli negatif, dan 0 dilabeli netral. Meskipun pelabelan berjalan otomatis, dataset luaran (output dataset) telah divalidasi secara kualitatif oleh pakar bahasa (Fitry Tunjung Saputri, M.Pd.) dan dinyatakan selaras dengan konteks makna riil di lapangan tanpa perlu anotasi manual ulang.

Dari total 2.417 dokumen, sebanyak 93 ulasan berskor netral (0) dieliminasi. Eliminasi ini dilakukan karena ulasan tersebut umumnya hanya memuat informasi objektif tanpa polaritas evaluatif yang tegas, sehingga berisiko mendilusi pembobotan fitur pada tahap *clustering*.

Hasil penyaringan menyisakan 2.324 dokumen beropini tegas yang diteruskan ke tahap pemodelan topik. Dari jumlah tersebut, distribusi komputasi menunjukkan bahwa mayoritas ulasan pasien bersifat positif, yakni sebanyak 1.937 ulasan (83,35%), sedangkan sisanya merupakan sentimen negatif sebanyak 387 ulasan (16,65%). Keberadaan sentimen negatif ini tetap signifikan dan menjadi dasar penting untuk analisis lanjutan pada tahap *data mining*. Secara visual, distribusi sentimen ulasan pasien disajikan pada Gambar 3.

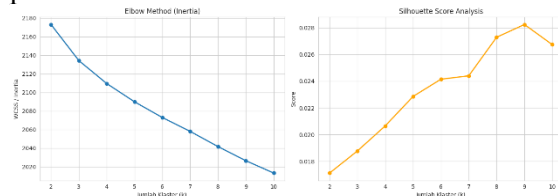
**Gambar 3** Distribusi Sentimen Ulasan Pasien

4.4. Transformation (TF-IDF)

Setelah dilakukan penyaringan, sebanyak 2.324 dokumen berpolaritas positif dan negatif ditransformasikan menjadi representasi vektor numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Transformasi ini bertujuan mengubah korpus tekstual menjadi matriks matematis agar jarak antar-dokumen dapat dihitung pada tahap clustering. Implementasi pembobotan menggunakan pustaka *TfidfVectorizer* dengan pengaturan *max_features = 1000* untuk membatasi dimensi matriks hanya pada 1.000 term teratas yang memiliki signifikansi tertinggi. Pembatasan ini secara efektif mereduksi noise dari kata langka sekaligus mempertahankan fitur yang paling representatif, sehingga menghasilkan matriks akhir berdimensi 2.324 baris (dokumen) dan 1.000 kolom (fitur kata) yang siap diproses pada algoritma data mining.

4.5. Data Mining (K-Means Clustering)

Matriks hasil pembobotan TF-IDF yang berdimensi (2324, 1000) diproses menggunakan algoritma *K-Means Clustering* untuk mengelompokkan ulasan berdasarkan kemiripan pola kata. Evaluasi penentuan jumlah kluster optimal (k) dilakukan pada rentang $k = 2$ hingga 10 menggunakan *Elbow Method* dan *Silhouette Score*, sebagaimana divisualisasikan pada Gambar 4.



Gambar 4 Hasil Evaluasi Jumlah Kluster: Elbow & Silhouette Score

Hasil perhitungan *Within-Cluster Sum of Squares* (WCSS) menunjukkan titik pelandaian (*elbow point*) yang paling optimal pada nilai $k = 4$, di mana penambahan kluster selanjutnya tidak lagi menghasilkan penurunan variasi intra-kluster yang signifikan. Meskipun validasi menggunakan *Silhouette Score* menghasilkan nilai absolut yang relatif rendah (0,0207 pada $k = 4$), kondisi batas kluster yang saling tumpang tindih (*overlapping*) ini merupakan fenomena lazim. Hal tersebut diakibatkan oleh *Curse of Dimensionality* pada metrik jarak matriks teks

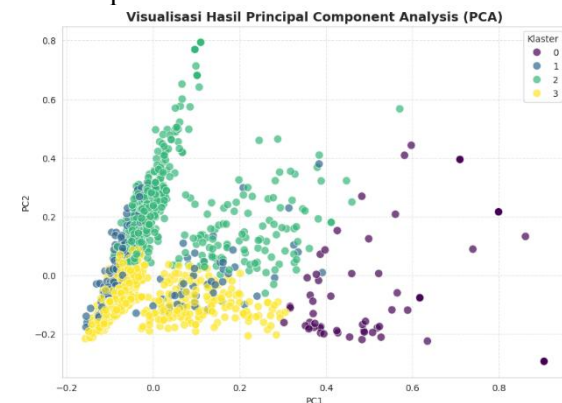
berdimensi sangat tinggi, serta karakteristik alamiah ulasan pasien yang kerap memuat keluhan multi-topik dalam satu kalimat. Oleh karena itu, $k = 4$ ditetapkan sebagai jumlah optimal karena memberikan keseimbangan terbaik antara kompleksitas komputasi dan interpretabilitas semantik (*actionable insights*).

Eksekusi pemodelan K-Means ($n_clusters = 4$, $random_state = 42$) mempartisi 2.324 dokumen ke dalam empat kluster yang stabil tanpa ada kluster kosong. Hasil pengelompokan menunjukkan distribusi yang tidak merata: Kluster 3 sangat mendominasi dengan 1.402 ulasan (60,33%), diikuti oleh Kluster 2 sebanyak 571 ulasan (24,57%), Kluster 1 sebanyak 245 ulasan (10,54%), dan Kluster 0 sebanyak 106 ulasan (4,56%). Disparitas proporsi ini mengindikasikan adanya satu area layanan spesifik yang menjadi pusat perhatian utama pasien, yang selanjutnya dibedah pada tahap interpretasi.

4.6. Evaluation & Interpretation

4.6.1. Validasi Visual Menggunakan PCA

Struktur kluster dievaluasi secara visual menggunakan teknik *Principal Component Analysis* (PCA). Data ulasan yang merepresentasikan matriks TF-IDF berdimensi tinggi (1.000 fitur) direduksi menjadi dua komponen utama (PC1 dan PC2). Visualisasi pada Gambar 5 menunjukkan bahwa algoritma K-Means berhasil memisahkan data ke dalam empat wilayah persebaran pola yang berbeda, meskipun terdapat tumpang tindih (*overlap*) sebagai konsekuensi logis dari karakteristik ulasan pasien yang kerap membahas keluhan lintas topik secara bersamaan.



Gambar 5 Visualisasi Reduksi Dimensi Kluster Menggunakan PCA

4.6.2. Interpretasi Topik dan Pelabelan Klaster

Interpretasi tematik dilakukan dengan mengekstraksi kata kunci berbobot TF-IDF tertinggi pada *centroid* setiap klaster, yang kemudian divalidasi secara kualitatif oleh pakar bahasa (Fitry Tunjung Saputri, M.Pd.). Berdasarkan analisis terhadap kata-kata dominan tersebut, makna tematik dari keempat klaster diuraikan sebagai berikut:

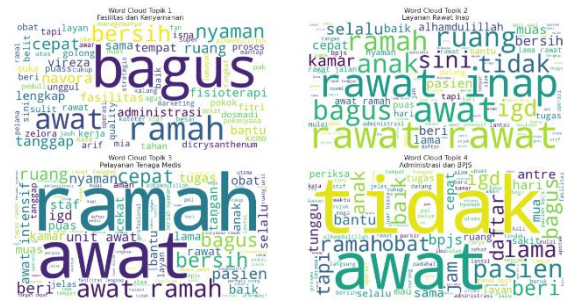
1. Klaster 0 dilabeli sebagai Fasilitas dan Kenyamanan karena didominasi kata bersih, nyaman, bagus, lengkap, dan navora (merujuk pada unit spesifik).
2. Klaster 1 diklasifikasikan sebagai Layanan Rawat Inap berkat kemunculan konsisten kata rawat, inap, ruang, dan kamar.
3. Klaster 2 menonjolkan kata ramah, tugas (petugas), tanggap, dan tangan, sehingga diinterpretasikan sebagai Pelayanan Tenaga Medis dengan fokus pada interaksi dan profesionalisme.
4. Klaster 3 dilabeli sebagai Administrasi dan BPJS akibat tingginya frekuensi kata bpjs, daftar, jam, lama, serta kata negasi (tidak). Hal ini mengindikasikan kuatnya sorotan pasien terhadap prosedur birokrasi.

Hasil penamaan topik beserta proporsinya dirangkum pada Tabel 4.

Tabel 4 Interpretasi Topik Dan Distribusi Ulasan Per Klaster

Klaster	Label Topik Layanan	Jumlah Ulasan	Proporsi
0	Fasilitas dan Kenyamanan	106	4,56%
1	Layanan Rawat Inap	245	10,54%
2	Pelayanan Tenaga Medis	571	24,57%
3	Administrasi dan BPJS	1402	60,33%

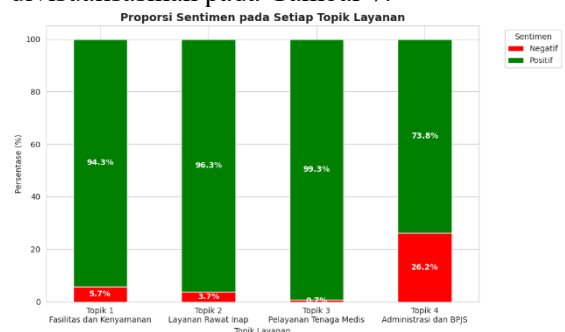
Guna memvalidasi representasi leksikal dari setiap penamaan klaster secara visual, dilakukan ekstraksi *Word Cloud* berdasarkan topik layanan sebagaimana disajikan pada Gambar 6. Pola visual tersebut konsisten dengan hasil pelabelan; Klaster Fasilitas dan Kenyamanan menonjolkan ukuran kata terkait lingkungan fisik, Klaster Rawat Inap memperlihatkan penekanan istilah ruang perawatan, Klaster Medis menonjolkan interaksi interpersonal, dan Klaster Administrasi didominasi oleh istilah proses pendaftaran dan antrian.



Gambar 6 Visualisasi *Word Cloud* Berdasarkan Topik Layanan

4.6.3. Integrasi Clustering dengan Analisis Sentimen

Untuk memperoleh gambaran komprehensif sekaligus mengevaluasi kinerja rumah sakit, hasil pengelompokan topik diintegrasikan dengan polaritas sentimen melalui teknik tabulasi silang (*cross-tabulation*). Pemetaan distribusi kepuasan pasien pada setiap klaster divisualisasikan pada Gambar 7.



Gambar 7 Proporsi Sentimen Pada Setiap Topik Layanan

Berdasarkan Gambar 7, integrasi ini mengungkap disparitas kualitas yang mencolok antar-area layanan Rumah Sakit XYZ. Topik Pelayanan Tenaga Medis teridentifikasi sebagai kekuatan utama dengan tingkat sentimen positif tertinggi mencapai 99,3%, disusul oleh Layanan Rawat Inap (96,3%) serta Fasilitas dan Kenyamanan (94,3%). Hal ini menunjukkan bahwa kompetensi klinis, keramahan tenaga kesehatan, serta kenyamanan lingkungan fisik telah memenuhi ekspektasi pasien secara optimal.

Sebaliknya, temuan paling krusial terdapat pada topik Administrasi dan BPJS yang merupakan titik kritis layanan. Topik ini sangat mendominasi diskusi pasien (60,33% dari total ulasan) dan memiliki rasio sentimen negatif tertinggi sebesar 26,2%. Sebagai langkah validasi akhir terhadap karakteristik kosakata

yang memicu kepuasan maupun keluhan tersebut, dilakukan ekstraksi *Word Cloud* sentimen pada Gambar 8.



Gambar 8. Visualisasi *Word Cloud* Berdasarkan Polaritas Sentimen

Visualisasi ketidakpuasan pada Gambar 8 secara mencolok menampilkan kata tidak, lama, tunggu, daftar, bpjs, dan obat. Kehadiran istilah tersebut memvalidasi secara leksikal bahwa sumber utama ketidakpuasan pasien bukan berasal dari aspek klinis, melainkan inefisiensi operasional dan kompleksitas sistem birokrasi.

4.6.4. Pembahasan (Implikasi Manajerial)

Secara umum, penilaian publik terhadap Rumah Sakit XYZ sangat baik (83,35% positif). Namun, temuan ini selaras dengan studi terdahulu yang menekankan bahwa penyederhanaan prosedur administratif BPJS merupakan determinan krusial dalam menentukan tingkat kepuasan pasien [26], [27], [28].

Secara metodologis, integrasi analisis sentimen dan pemodelan topik pada penelitian ini terbukti efektif mengekstrak informasi strategis (*actionable insights*). Sebagai implikasi manajerial, prioritas perbaikan layanan rumah sakit harus difokuskan pada pembenahan sistem administratif. Solusi fungsional yang direkomendasikan adalah implementasi sistem informasi yang terintegrasi secara komprehensif, seperti peningkatan interoperabilitas antara Sistem Informasi Manajemen Rumah Sakit (SIMRS) dengan aplikasi Mobile JKN, serta adopsi layanan resep elektronik (*e-prescription*) guna mereduksi waktu tunggu pasien secara terukur.

5. KESIMPULAN

Berdasarkan hasil analisis *text mining* terhadap ulasan pasien Rumah Sakit XYZ, dapat ditarik beberapa kesimpulan sebagai berikut:

- Penerapan algoritma *K-Means Clustering* berbasis representasi TF-IDF terbukti efektif mengekstraksi struktur topik dari teks tidak terstruktur. Model berhasil

memetakan 2.324 ulasan ke dalam empat area layanan utama: Administrasi dan BPJS (60,33%), Pelayanan Tenaga Medis (24,57%), Layanan Rawat Inap (10,54%), serta Fasilitas dan Kenyamanan (4,56%).

- Pendekatan *Lexicon-Based Scoring* dengan fitur *negation handling* berhasil melabeli persepsi publik secara akurat tanpa data latih. Hasilnya menunjukkan dominasi sentimen positif sebesar 83,35% dan negatif sebesar 16,65%. Pada tingkat klaster, kepuasan tertinggi dicapai pada Pelayanan Tenaga Medis (99,3%), sedangkan tingkat ketidakpuasan terpusat pada Administrasi dan BPJS (26,2%).
- Kelebihan dari integrasi metode klusterisasi dan leksikon ini adalah kemampuannya menghasilkan wawasan berbasis data (*actionable insights*) yang terarah. Secara manajerial, pemetaan ini merekomendasikan manajemen rumah sakit untuk memprioritaskan perbaikan operasional dan efisiensi waktu tunggu pada ranah Administrasi dan BPJS, misalnya melalui peningkatan interoperabilitas SIMRS dengan aplikasi Mobile JKN.
- Untuk kemungkinan pengembangan selanjutnya, disarankan menggunakan pendekatan *Deep Learning* (seperti IndoBERT) guna mengatasi keterbatasan leksikon dalam memahami semantik kompleks. Peneliti juga dapat membandingkan metode pemodelan dengan algoritma probabilistik seperti *Latent Dirichlet Allocation* (LDA) untuk menangani ulasan multitema, serta memperluas akuisisi data (*multichannel scraping*) untuk dibangun menjadi sebuah *dashboard* analitik *real-time*.

UCAPAN TERIMA KASIH

Penulis bersyukur kepada Allah SWT atas selesainya penelitian ini. Terima kasih sebesar-besarnya disampaikan kepada Bapak Ade Andri Hendriadi, M.Kom. dan Bapak Ahmad Khusaeri, M.Kom. atas segala bimbingan dan arahannya. Penulis juga mengucapkan terima kasih yang mendalam kepada kedua orang tua dan seseorang yang "istimewa" atas doa serta dukungan moral yang tidak pernah putus.

Terakhir, apresiasi ditujukan untuk sahabat-sahabat "6s" (Galih, Rangga, Eki, Deft, dan Gerald) yang selalu menjadi *support system* terbaik dari awal hingga akhir penelitian.

DAFTAR PUSTAKA

- [1] A. C. T. Angel, V. H. Pranatawijaya, and Widiatry, "Analisis Sentimen dan Emosi dari Ulasan Google Maps untuk Layanan Rumah Sakit di Palangka Raya Menggunakan Machine Learning," *KONSTELASI: Konvergensi Teknologi dan Sistem Informasi*, vol. 4, no. 1, Jun. 2024.
- [2] P. G. G. Wahyu, I. Setiawan, and R. I. Saputri, "Perilaku Pencarian Informasi Kesehatan Melalui Internet di Masyarakat," *Padjadjaran Journal of Dental Researchers and Students*, vol. 7, p. 81, Feb. 2023, doi: 10.24198/pjdrs.v6i2.40474.
- [3] H. Purba and I. Irwansyah, "User Generated Content dan Pemanfaatan Media Sosial Dalam Perkembangan Industri Pariwisata: Literature Review," *Jurnal Professional: Jurnal Komunikasi & Administrasi Publik*, vol. 9, no. 2, pp. 229–238, Dec. 2022, [Online]. Available: <https://www.researchgate.net/publication/371539936>
- [4] F. I. Syaputra and S. Wijayanto, "Analisis Sentimen Layanan Rumah Sakit di Purwokerto Berdasarkan Ulasan Google Maps Menggunakan Metode LSTM," *e-Proceeding of Engineering*, vol. 12, no. 4, pp. 5763–5768, Aug. 2025.
- [5] M. M. Bujnowska-Fedak and P. Węgierek, "The impact of online health information on patient health behaviours and making decisions concerning health," *Int. J. Environ. Res. Public Health*, vol. 17, no. 3, Jan. 2020, doi: 10.3390/ijerph17030880.
- [6] Y. Zhu and C. W. Piercy, "Impact of Online Patient Reviews: Three Choice-based Conjoint Experiments of Hospital and Physician Ratings," *Journal of Communication Technology*, vol. 6, no. 1, Apr. 2024, doi: 10.51548/joctec.6.1.2024.01.
- [7] A. F. Rizki, W. Prihartono, and Fathurrohman, "Analisis Sentimen Ulasan Google Maps Rumah Sakit Khalishah Di Cirebon Dengan Algoritma Naive Bayes," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 2, Apr. 2025, doi: 10.23960/jitet.v13i2.6309.
- [8] N. S. W. Al Qorni, "Penerapan Algoritma K-Means Clustering Pada Ulasan Pengguna Merdeka Mengajar Di Play Store," Aug. 2024.
- [9] I. Nurhakim and A. Voutama, "Analisis Efisiensi Pelayanan Kesehatan Dengan Visualisasi Data Interaktif Di Power BI," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 2, Apr. 2025, doi: 10.23960/jitet.v13i2.6355.
- [10] D. Musfiroh, U. Khaira, P. E. P. Utomo, and T. Suratno, "Analisis Sentimen terhadap Perkuliahan Daring di Indonesia dari Twitter Dataset Menggunakan InSet Lexicon," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 1, no. 1, pp. 24–33, Apr. 2021.
- [11] R. F. Ramadhan, S. H. Wijoyo, and M. C. Saputra, "Penerapan Metode K-Means Clustering pada Ulasan Perumahan PT XYZ di Google Maps untuk Formulasi Strategi Bisnis dengan Analisis SWOT," *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, vol. 7, no. 6, pp. 2879–2888, Jun. 2023, [Online]. Available: <http://j-ptiik.ub.ac.id>
- [12] E. Kurniawan and N. Hendrastuty, "Penerapan Algoritma K-Means Untuk Melakukan Klasterisasi Pada Peringkasan Teks," *Jurnal Informatika Teknologi dan Sains (JINTEKS)*, vol. 6, Aug. 2024.
- [13] A. E. Widjaja, A. Fransisko, C. A. Haryani, and Hery, "Text Mining Application with K-Means Clustering to Identify Sentiments and Popular Topics: a Case Study of the three Largest Online Marketplaces in Indonesia," *Journal of Applied Data Sciences*, vol. 4, no. 4, pp. 441–453, Dec. 2023, doi: 10.47738/jads.v4i4.134.
- [14] A. S. R. Rufaida, A. E. Permanasari, and N. A. Setiawan, "Lexicon-Based Sentiment Analysis Using Inset Dictionary: A Systematic Literature Review," *European Alliance for Innovation n.o.*, Jun. 2023. doi: 10.4108/eai.5-10-2022.2327474.
- [15] W. O. Puspitasari, M. Mustamin, J. Fitrianiingsih, and M. T. Malik, "Analisis Status Akreditasi Rumah Sakit terhadap Kepuasan Tenaga Kesehatan dan Kepuasan Pasien di Rumah Sakit Hikmah Masamba," *Syntax Admiration*, vol. 5, no. 12, Dec. 2024.
- [16] T. Jo, *Text Mining Concepts, Implementation, and Big Data Challenge*, vol. 45. Springer International Publishing AG, part of Springer Nature, 2019. [Online]. Available: <http://www.springer.com/series/11970>
- [17] P. Permatasari and N. A. Rajebta, "Improving Quality of Care on Patient Satisfaction in Health Service Facilities by Rating on Google Maps: Literature Review," *Care: Jurnal Ilmiah Ilmu Kesehatan*, vol. 13, no. 2, pp. 276–285, Jul. 2025, doi: 10.33366/jc.v13i2.6454.
- [18] A. H. Dani, E. Y. Puspaningrum, and R. Mumpuni, "Studi Performa TF-IDF dan

- Word2Vec Pada Analisis Sentimen Cyberbullying,” vol. 2, no. 2, pp. 94–106, Jun. 2024, doi: 10.62951/router.v2i2.76.
- [19] R. Ramadhan, Y. A. Sari, and P. P. Adikara, “Perbandingan Pembobotan Term Frequency-Inverse Document Frequency dan Term Frequency-Relevance Frequency terhadap Fitur N-Gram pada Analisis Sentimen,” Nov. 2021. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [20] J. Han, M. Kamber, and J. Pei, “Data Mining. Concepts and Techniques, 3rd Edition (The Morgan Kaufmann Series in Data Management Systems),” 2011.
- [21] P. Vania and B. Nurina Sari, “Perbandingan Metode Elbow dan Silhouette untuk Penentuan Jumlah Klaster yang Optimal pada Clustering Produksi Padi menggunakan Algoritma K-Means,” *Jurnal Ilmiah Wahana Pendidikan*, vol. 9, no. 21, pp. 547–558, 2023, doi: 10.5281/zenodo.10081332.
- [22] S. A. Nugraha, “Penerapan Lexicon Based Untuk Analisis Sentimen Masyarakat Indonesia Terhadap Danantara,” *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 3, Jun. 2025.
- [23] F. Koto and G. Y. Rahmaningtyas, “InSet Lexicon: Evaluation of a Word List for Indonesian Sentiment Analysis in Microblogs,” *2017 International Conference on Asian Language Processing (IALP)*, pp. 391–394, Jul. 2017, doi: 10.1109/IALP.2017.8300625.
- [24] T. Eng, M. R. Ibn Nawab, and K. M. Shahiduzzaman, “Improving Accuracy of The Sentence-Level Lexicon-Based Sentiment Analysis Using Machine Learning,” *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, pp. 57–69, Jan. 2021, doi: 10.32628/cseit21717.
- [25] N. A. Salsabila, “Kamus Alay - Colloquial Indonesian Lexicon,” GitHub. Accessed: Dec. 09, 2025. [Online]. Available: <https://github.com/nasalsabila/kamus-alay>
- [26] S. Sutarji, R. Wardani, A. Bastian, and P. Radono, “Journal of Hospital Management and Services The Relationship Of The Quality Of Administrative Services And Bpjs Patient Satisfaction In The Inpatient Room X Rumkit Tk Ii Dr. Soepraoen Kesdam V/Brawijaya Malang In Malang City,” *Journal of Hospital Management and Services*, vol. 5, no. 2, pp. 55–60, Nov. 2023, doi: <https://doi.org/10.30994/jhms.v5i2.56>.
- [27] R. A. Agnaty, A. Andriyani, and S. Jaksa, “Analisis Kepuasan Pasien Bpjs Terhadap Kualitas Pelayanan Rawat Jalan Di Rumah Sakit Dengan Pendekatan Servqual,” *Jurnal Manajemen Pelayanan Kesehatan*, vol. 28, no. 2, pp. 44–49, Jun. 2025, doi: <https://doi.org/10.22146/jmpk.v28i02.21340>.
- [28] D. Syahfitri, D. Rizka, A. Pulungan, and F. P. Gurning, “Analysis of the Level of Service Satisfaction of BPJS Patients at the Primary Outpatient Clinic of UIN North Sumatra Using the Servqual Method,” *Jurnal Kesmas Prima Indonesia (JKPI)*, vol. 9, no. 1, Jul. 2025, doi: <https://doi.org/10.34012/jkpi.v9i2.7210>.