

TUGAS AKHIR

ANALISIS SENTIMEN BERBASIS ASPEK PADA PELAYANAN HOTEL DI PEKANBARU MENGGUNAKAN METODE NAÏVE BAYES DAN LOGISTIC REGRESSION

UNIVERSITAS ISLAM RIAU



FADHIA HAYA
203510487

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS TEKNIK
UNIVERSITAS ISLAM RIAU
PEKANBARU
2026**

HALAMAN PENGESAHAN TUGAS AKHIR

Nama : Fadhia Haya
NPM : 203510487
Kelompok Keahlian : Artificial Intelligence
Program Studi : Teknik Informatika
Jenjang Pendidikan : Strata Satu (S1)
Judul TA : Analisis Sentimen Berbasis Aspek Pada Pelayanan Hotel
Di Pekanbaru Menggunakan Metode Naïve Bayes dan
Logistic Regression

Format sistematika dan pembahasan materi pada masing-masing bab dan sub bab dalam tugas akhir ini telah dipelajari dan dinilai relatif telah memenuhi ketentuan-ketentuan dan kriteria-kriteria dalam metode penelitian ilmiah. Oleh karena itu tugas akhir ini dinilai layak dapat disetujui untuk disidangkan dalam ujian Seminar Tugas Akhir.

Pekanbaru, 26 Juni 2025

Di sahkan oleh :

Penguji I

Penguji II


Dr. Arbi Haza Nasution, B.IT., M.IT.


Rizky Wandri, S.Kom., M.Kom.

Ketua Program Studi
Teknik Informatika

Dosen Pembimbing


Dr. Ir. Des Suryani, M.Sc.


Anggi Hanafiah, S.Kom., M.Kom.

HALAMAN PENGESAHAN DEWAN PENGUJI TUGAS AKHIR

Nama : Fadhia Haya
NPM : 203510487
Kelompok Keahlian : Artificial Intelligence
Program Studi : Teknik Informatika
Jenjang Pendidikan : Strata Satu (S1)
Judul TA : Analisis Sentimen Berbasis Aspek Pada Pelayanan Hotel di Pekanbaru Menggunakan Metode Naive Bayes dan Logistic Regression

Tugas Akhir ini secara keseluruhan dinilai telah memenuhi ketentuan-ketentuan dan kaidah-kaidah dalam penulisan penelitian ilmiah serta telah diuji dan dapat dipertahankan dihadapan dewan penguji. Oleh karena itu, Tim Penguji Ujian Tugas Akhir Fakultas Teknik Universitas Islam Riau menyatakan bahwa mahasiswa yang bersangkutan dinyatakan Telah Lulus Mengikuti Ujian Tugas Akhir Pada Tanggal 15 Januari 2026 dan disetujui serta diterima untuk memenuhi salah satu syarat guna memperoleh gelar Sarjana Strata Satu Bidang Ilmu Teknik Informatika.

Pekanbaru, 15 Januari 2026

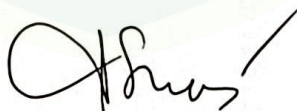
Dewan Penguji

1. Pembimbing : Anggi Hanafiah, S.Kom., M.Kom.
2. Penguji 1 : Dr. Arbi Haza Nasution B.IT. (Hons). M.IT.
3. Penguji 2 : Rizky Wandri, S.Kom., M.Kom

()
()

Disahkan Oleh :

Ketua Program Studi
Teknik Informatika



Dr. Ir. Des Suryani, M.Sc

NIDN. 1026126801

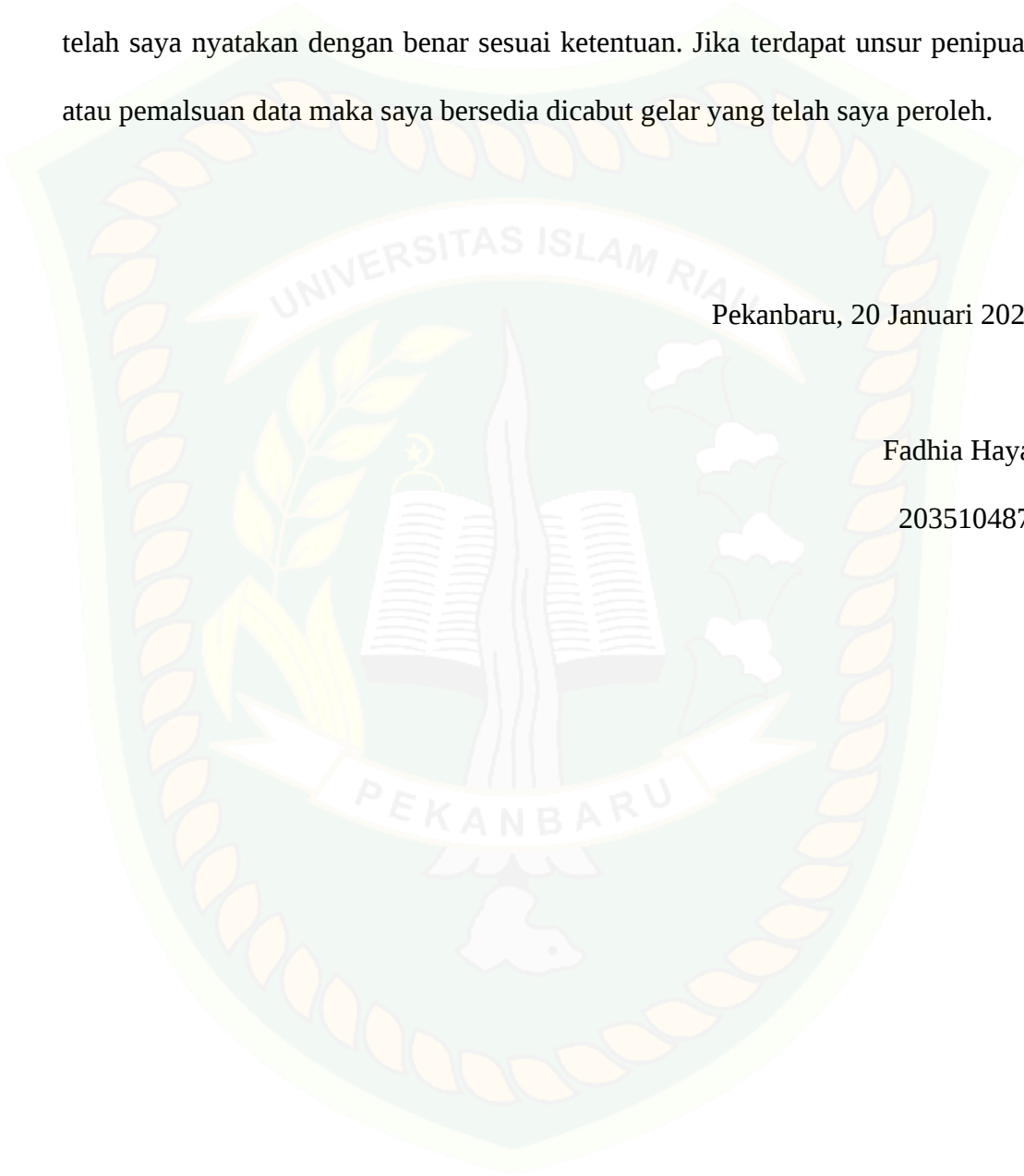
PERNYATAAN KEASLIAN TUGAS AKHIR

Dengan ini saya menyatakan bahwa tugas akhir ini merupakan karya saya sendiri dan semua sumber yang tercantum didalamnya baik yang dikutip maupun dirujuk telah saya nyatakan dengan benar sesuai ketentuan. Jika terdapat unsur penipuan atau pemalsuan data maka saya bersedia dicabut gelar yang telah saya peroleh.

Pekanbaru, 20 Januari 2026

Fadhia Haya

203510487



KATA PENGANTAR

Assalamu'alaikum Warahmatullahi Wabarakatuh.

Dengan mengucapkan Alhamdulillahirobbil'alamin, berkat rahmat dan hidayah Allah SWT serta nikmat yang tak terhingga, penulis dapat menyelesaikan proposal skripsi ini dengan judul "Analisis Sentimen Berbasis Aspek pada Pelayanan Hotel di Pekanbaru Menggunakan Metode *Naïve Bayes* dan Logistic Regression" sebagai salah satu syarat wajib untuk mendapatkan gelar sarjana pada Fakultas Teknik Program Studi Teknik Informatika Universitas Islam Riau. Dalam proses pembuatan proposal skripsi ini, penulis menyadari bahwa tanpa bantuan dan bimbingan sebagai pihak maka proposal ini sulit untuk terwujud. Untuk itu dalam kesempatan ini penulis ingin menyampaikan ucapan terima kasih yang sebesar-besarnya kepada:

1. Ir. Dr. Des Suryani., M.Sc selaku ketua Program Studi Teknik Informatika Universitas Islam Riau.
2. Bapak Anggi Hanafiah., S.Kom., M.Kom selaku sekretaris Program Studi Teknik Informatika dan dosen pembimbing yang telah memberikan ilmu, motivasi, nasihat yang bermanfaat, serta ikhlas dan sabar memberikan bimbingan dan arahan disela-sela kesibukan beliau kepada penulis untuk menyelesaikan tugas akhir ini.
3. Kepada seluruh dosen Teknik Informatika yang telah membimbing dan membantu penulis pada proses belajar mengajar selama di bangku perkuliahan.
4. Kedua orang tua paling berjasa dalam hidup penulis, Papah Silvia dan Mamah Ludia Nata. Terima kasih atas kepercayaan yang telah diberikan kepada penulis untuk melanjutkan Pendidikan kuliah, serta cinta, do'a, motivasi, semangat dan nasihat yang tiada hentinya diberikan kepada penulis dalam penyusunan skripsi.

5. Kepada saudara satu-satunya penulis, Bharada M.Alfauzan yang telah memberikan semangat, motivasi dan menyakinkan penulis bisa untuk menyelesaikan skripsi ini.
6. Untuk seseorang yang telah membersamai penulis yang tidak bisa disebutkan namanya, terima kasih atas dukungan, motivasi, do'a serta cinta yang telah diberikan kepada penulis, serta terima kasih setia meluangkan waktunya untuk menjadi tempat dan pendengar terbaik penulis sampai akhirnya penulis dapat menyelesaikan skripsi ini.
7. Kepada seluruh teman penulis yang senantiasa memberikan dukungan untuk penulis dalam penyelesaian skripsi ini hingga selesai.
8. Kepada diri saya sendiri Fadhia Haya, apresiasi yang sebesar-besarnya terima kasih telah bertanggung jawab untuk menyelesaikan apa yang telah di mulai. Terima kasih untuk terus berusaha dan tidak menyerah, serta menikmati setiap prosesnya yang bisa dibilang tidak mudah.

Penulis menyadari masih banyak kekurangan dalam penyusunan laporan skripsi ini, untuk itu dengan segala kerendahan hati penulis mengharapkan saran dan kritik yang sifatnya membangun dari pembaca untuk penyempurnaan laporan skripsi ini. Akhir kata, semoga laporan skripsi ini dapat menambah pengetahuan dan bermanfaat bagi para pembaca.

Wassalamu'alaikum Warahmatullahi Wabarakatuh.

Pekanbaru, 1 Oktober 2024

Fadhia Haya

DAFTAR ISI

KATA PENGANTAR.....	i
DAFTAR ISI.....	iii
DAFTAR GAMBAR.....	vi
DAFTAR TABEL.....	vii
ABSTRAK	viii
ABSTRACT	ix
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang.....	1
1.2 Identifikasi Masalah	3
1.3 Rumusan Masalah	4
1.4 Batasan Masalah.....	5
1.5 Tujuan Penelitian.....	5
1.6 Manfaat Penelitian.....	6
BAB II LANDASAN TEORI	7
2.1 Tinjauan Pustaka	7
2.2 Dasar Teori	13
2.2.1 <i>Text Mining</i>	13
2.2.1.1 Analisis Sentimen.....	13
2.2.1.2 Pemodelan Topik.....	14
2.2.1.3 <i>Text Generation</i>	14
2.2.2 Ulasan	15
2.2.3 Google Maps.....	15
2.2.4 Web Scrapping	16
2.2.5 Pra-pemrosesan Data	17
2.2.6 Teknik <i>Prompting</i>	17
2.2.6.1 Zero-Shot.....	18
2.2.6.2 Few-Shot	18
2.2.7 Imbalanced Data	19
2.2.8 Ekstraksi Fitur.....	19
2.2.8.1 N-Gram	20

2.2.8.2 TF-IDF	20
2.2.9 Algoritma <i>Text Mining</i>	21
2.2.9.1 <i>Latent Dirichlet Allocation</i>	21
2.2.9.2 <i>Naïve Bayes</i>	22
2.2.9.3 <i>Logistic Regression</i>	23
2.2.10 <i>K-Fold Cross Validation</i>	24
2.2.11 Metrik Evaluasi	25
2.2.11.1 <i>Topic Coherence</i>	25
2.2.11.2 <i>Perplexity</i>	26
2.2.12 <i>Elbow Method</i>	27
2.2.12.1 Akurasi	27
2.3 Kerangka Pemikiran	28
BAB III METODOLOGI PENELITIAN	30
3.1 Tempat Penelitian	30
3.2 Alat dan Bahan	30
3.3 Tahapan Penelitian	31
3.3.1 Pengumpulan Data	32
3.3.2 Pra-pemrosesan Data	32
3.3.2.1 <i>Case Folding</i>	32
3.3.2.2 Penghapusan Emoji, Angka, dan Tanda Baca	33
3.3.2.3 Penghapusan <i>Whitespace</i> dan Karakter Tunggal	33
3.3.2.4 Penghapusan Kata Tunggal	34
3.3.2.5 Penghapusan Kalimat Berbahasa Inggris	34
3.3.2.6 Koreksi Ejaan dan Singkatan	34
3.3.2.7 Tokenisasi	35
3.3.2.8 Penghapusan <i>Stopwords</i>	35
3.3.2.9 <i>Stemming</i>	35
3.3.3 Anotasi Data	36
3.3.3.1 <i>Zero-Shot Prompting</i>	36
3.3.4 Pemodelan Topik	37
3.3.5 Evaluasi Pemodelan Topik	37
3.3.6 K-Fold Cross Validation	38

3.3.7 Ekstraksi Fitur.....	38
3.3.8 SMOTE.....	40
3.3.9 Analisis Sentimen Berbasis Aspek.....	40
3.3.9.1 <i>Naïve Bayes</i>	40
3.3.9.2 <i>Logisitic Regression</i>	42
3.3.10 Evaluasi.....	44
BAB IV HASIL DAN PEMBAHASAN.....	45
4.1 Pengumpulan Data.....	45
4.2 Pra-pemrosesan Data.....	45
4.3 Anotasi Data.....	47
4.4 Pemodelan Topik.....	48
4.5 <i>Elbow Method</i>	48
4.5.1 <i>Coherence Score</i>	49
4.5.2 <i>Perplexity</i>	50
4.6 Penentuan Jumlah Topik.....	50
4.7 Analisis Sentimen Berbasis Aspek.....	53
4.8 Pengujian Model Klasifikasi.....	66
BAB V KESIMPULAN.....	68
5.1 Kesimpulan.....	68
5.2 Saran.....	69
DAFTAR PUSTAKA.....	70

DAFTAR GAMBAR

Gambar 2.1 Notasi Grafis Model LDA.....	22
Gambar 2.2 Alur Kerja <i>K-Fold Cross Validation</i>	25
Gambar 2.3 Proses Perhitungan <i>Coherence Score</i>	26
Gambar 2.4 Kerangka Pemikiran.....	29
Gambar 3.1 Gambaran Umum Penelitian.....	31
Gambar 4.1 Hasil Pengumpulan Data.....	45
Gambar 4.2 Persentase Sentimen.....	48
Gambar 4.3 Perbandingan <i>Coherence Value</i>	49
Gambar 4.4 Perbandingan <i>Perplexity Value</i>	50
Gambar 4.5 Visualisasi LDA	51
Gambar 4.6 Perbandingan Akurasi Tiap <i>Fold</i> Klasifikasi Aspek	55
Gambar 4.7 <i>Confussion Matrix</i> Model <i>Naïve Bayes</i> Pada Klasifikasi Aspek.....	56
Gambar 4.8 <i>Confussion Matrix</i> Model <i>Logistic Regression</i> Pada Klasifikasi Aspek	56
Gambar 4.9 Perbandingan Akurasi Tiap <i>Fold</i> Klasifikasi Sentimen Aspek Fasilitas	58
Gambar 4.10 <i>Confusion Matrix</i> Model <i>Naïve Bayes</i> pada Klasifikasi Sentimen Aspek Fasilitas	59
Gambar 4.11 <i>Confusion Matrix</i> Model <i>Logistic Regression</i> pada Klasifikasi Sentimen Aspek Fasilitas	59
Gambar 4.12 Perbandingan Akurasi Tiap <i>Fold</i> Klasifikasi Sentimen Aspek Pelayanan	60
Gambar 4.13 <i>Confusion Matrix</i> Model <i>Naïve Bayes</i> pada Klasifikasi Sentimen Aspek Pelayanan	61
Gambar 4.14 <i>Confusion Matrix</i> Model <i>Logistic Regression</i> pada Klasifikasi Sentimen Aspek Pelayanan	62
Gambar 4.15 Perbandingan Akurasi Tiap <i>Fold</i> Klasifikasi Sentimen Aspek Lokasi	63
Gambar 4.16 <i>Confusion Matrix</i> Model <i>Naïve Bayes</i> pada Klasifikasi Sentimen Aspek Lokasi.....	64
Gambar 4.17 <i>Confusion Matrix</i> Model <i>Logistic Regression</i> pada Klasifikasi Sentimen Aspek Lokasi.....	65

DAFTAR TABEL

Tabel 2.1 Rangkuman Penelitian Terdahulu	10
Tabel 3.1 <i>Case Folding</i>	33
Tabel 3.2 Penghapusan Emoji, Angka, dan Tanda Baca	33
Tabel 3.3 Penghapusan <i>Whitespace</i> dan Karakter Tunggal	33
Tabel 3.4 Koreksi Ejaan dan Singkatan	34
Tabel 3.5 Tokenisasi	35
Tabel 3.6 Penghapusan <i>Stopwords</i>	35
Tabel 3.7 <i>Stemming</i>	36
Tabel 3.8 N-Grams	38
Tabel 3.9 TF-IDF.....	39
Tabel 3.10 Transformasi Ekstraksi Fitur pada Data Uji	40
Tabel 3.11 Rata-rata dan Standar Deviasi Kelas A	41
Tabel 3.12 Rata-rata dan Standar Deviasi Kelas B	41
Tabel 3.13 Asumsi Nilai Koefisien	43
Tabel 4.1 Sebelum dan Sesudah Pra-pemrosesan	46
Tabel 4.2 Hasil Anotasi Sentimen	47
Tabel 4.3 Nilai <i>Coherence</i> dan <i>Perplexity</i>	49
Tabel 4.4 Kata Kunci Model Pemodelan Topik	52
Tabel 4.5 Distribusi Aspek	53
Tabel 4.6 Rata-rata Akurasi Model Klasifikasi Aspek	54
Tabel 4.7 Rata-rata Akurasi Model Klasifikasi Sentimen.....	57
Tabel 4.8 Data Pengujian	66
Tabel 4.9 Hasil Pengujian	67

ABSTRAK

Perkembangan industri perhotelan di Kota Pekanbaru mendorong peningkatan jumlah ulasan pelanggan pada platform daring, khususnya Google Maps. Ulasan tersebut memuat berbagai opini pelanggan terhadap aspek pelayanan hotel yang dapat dimanfaatkan sebagai bahan evaluasi kualitas layanan. Namun, besarnya volume data menjadikan analisis manual kurang efektif, sehingga diperlukan pendekatan otomatis melalui analisis sentimen berbasis aspek. Penelitian ini bertujuan untuk menganalisis sentimen berbasis aspek terhadap ulasan hotel di Pekanbaru serta membandingkan kinerja metode Naïve Bayes dan Logistic Regression. Data diperoleh melalui teknik web scraping, kemudian diproses menggunakan tahapan prapemrosesan teks dan ekstraksi fitur TF-IDF. Untuk mengatasi ketidakseimbangan data, diterapkan metode Synthetic Minority Over-sampling Technique (SMOTE). Proses klasifikasi dilakukan menggunakan algoritma Naïve Bayes dan Logistic Regression dengan skema evaluasi Stratified K-Fold Cross Validation sebanyak lima fold. Hasil penelitian menunjukkan bahwa metode Logistic Regression memberikan kinerja yang lebih baik dibandingkan Naïve Bayes, dengan rata-rata akurasi sebesar **91,16%**, sedangkan Naïve Bayes mencapai **87,38%**. Dengan demikian, Logistic Regression lebih efektif digunakan dalam analisis sentimen berbasis aspek pada ulasan pelanggan hotel di Pekanbaru. Hasil penelitian ini diharapkan dapat membantu pihak hotel dalam meningkatkan kualitas pelayanan serta menjadi referensi bagi penelitian selanjutnya.

Kata kunci: Analisis Sentimen, Berbasis Aspek, Naïve Bayes, Logistic Regression, Hotel, Google Maps.

ABSTRACT

The growth of the hospitality industry in Pekanbaru City has led to an increasing number of customer reviews on online platforms, particularly Google Maps. These reviews contain various customer opinions regarding hotel service aspects, which can be utilized as a basis for evaluating service quality. However, the large volume of data makes manual analysis inefficient, thus requiring an automated approach through aspect-based sentiment analysis. This study aims to perform aspect-based sentiment analysis on hotel reviews in Pekanbaru and to compare the performance of the Naïve Bayes and Logistic Regression methods. The data were collected using web scraping techniques and processed through text preprocessing and TF-IDF feature extraction. To address data imbalance, the Synthetic Minority Over-sampling Technique (SMOTE) was applied. The classification process was conducted using Naïve Bayes and Logistic Regression algorithms with a five-fold Stratified K-Fold Cross Validation evaluation scheme. The results indicate that Logistic Regression outperforms Naïve Bayes, achieving an average accuracy of **91.16%**, while Naïve Bayes attains **87.38%**. Therefore, Logistic Regression is more effective for aspect-based sentiment analysis of hotel customer reviews in Pekanbaru. The findings of this study are expected to assist hotel management in improving service quality and to serve as a reference for future research.

Keywords: Sentiment Analysis, Aspect-Based, Naïve Bayes, Logistic Regression, Hotel, Google Maps.

BAB I

PENDAHULUAN

1.1 Latar Belakang

Industri perhotelan di Kota Pekanbaru mengalami perkembangan pesat seiring meningkatnya aktivitas bisnis dan pariwisata. Dalam era digital, ulasan atau *review* dari pengunjung hotel yang dipublikasikan melalui berbagai platform *online*, seperti Google Maps, Traveloka, dan Agoda, menjadi sumber informasi penting bagi calon wisatawan dalam menentukan pilihan tempat menginap. *Review* hotel menjadi salah satu dataset yang sangat penting untuk diteliti karena *review* terkait karakteristik di setiap hotel selalu berbeda (Tran Xuan et al., 2019). *Review* tersebut menjadi sumber informasi bagi para wisatawan untuk memilih tempat singgah saat dalam perjalanan. Selain itu, *review* juga menjadi sarana pengembangan kualitas hotel oleh pemilik usaha. Oleh karena itu, data *review* hotel dapat diolah untuk mengetahui penilaian yang terkandung dalam *review* dan memberikan informasi pengembangan industri perhotelan. Hotel-hotel tersebut akan menyediakan berbagai pelayanan untuk para wisatawan yang menginap. Kualitas layanan hotel dapat dilihat dari opini-opini diberikan oleh pengunjung yang menginap di hotel (Setiowati & Helen, 2018). Para wisatawan akan menuliskan pengalaman tentang layanan yang dialaminya pada platform *online* seperti Google Maps. Ulasan-ulasan dari pengunjung hotel akan mempengaruhi keputusan dari calon pengunjung baru hotel tersebut (Amira et al., 2020).

Google Maps menyediakan informasi mengenai arah, lokasi, serta rute terbaik untuk mencapai tempat yang ingin dikunjungi. Ulasan dari pengguna dapat membentuk persepsi kita terhadap sejauh mana sebuah lokasi disukai atau tidak. Hal ini tentu menjadi bahan pertimbangan bagi pihak-pihak tertentu yang ingin memahami opini atau sentimen masyarakat terhadap suatu tempat. Namun, penulis menilai bahwa penilaian berupa *rating* dan komentar yang diberikan pengguna tidak selalu akurat, karena masih ditemukan ulasan dengan *rating* bintang 5 yang disertai komentar negatif atau tidak pantas. Hal ini tentu menjadi perhatian dan pertimbangan tersendiri bagi masyarakat (Erfina & Resti Wardani, 2022). Banyaknya ulasan yang diterima membuat analisis manual menjadi tidak praktis dan memakan waktu sehingga manajemen hotel sulit untuk mengambil sentimen pelanggan secara keseluruhan. Salah satu Teknik yang dapat dilakukan untuk menyelesaikan masalah ini adalah analisis sentimen berbasis aspek.

Pada penelitian sebelumnya (Gulati et al., 2022) membandingkan tujuh algoritma *machine learning* dengan implementasi N-Gram yang melakukan *sentiment analysis* terhadap 72.000 *tweets* di Twitter yang berhubungan dengan Covid-19. Terdapat tiga tipe N-Gram yang digunakan yakni *Unigram*, *Bigram*, dan *Trigram*. Algoritma *Logistic Regression* mendapatkan *average accuracy* sebesar 0.97. *Logistic Regression* adalah metode *supervised learning* untuk menganalisis data terhadap satu atau lebih variabel prediksi dengan satu variabel respon (Assaidi & Amin, 2022). Penelitian lebih lanjut yang dilakukan oleh (Helmayanti et al., 2023) menguatkan keefektifan metode *Naive Bayes* dalam melakukan analisis sentimen terhadap ulasan pengguna mengenai aplikasi pada platform Google Play

Store. Hasil penelitian menunjukkan bahwa *Naive Bayes* mampu mencapai tingkat akurasi yang baik dalam mengklasifikasikan teks yang berorientasi pada sentimen, menjadikan metodologi yang cocok untuk menganalisis ulasan layanan di hotel yang ada di Pekanbaru. Melalui penelitian ini penulis dapat mengetahui klasifikasi tanggapan pengunjung terhadap *review* pelayanan dan fasilitas hotel sebagai bahan masukan dan evaluasi bagi beberapa hotel yang ada di Pekanbaru dengan menerapkan metode *Naïve Bayes* dan *logistic regression*.

Analisis sentimen merupakan salah satu teknik yang dapat digunakan untuk menganalisis ulasan dari pengunjung hotel. Analisis sentimen dapat digunakan untuk melihat suatu opini yang ditujukan untuk hotel termasuk positif maupun negatif (Setiawan et al., 2022). Banyak penelitian yang melakukan analisis sentimen pada suatu ulasan, Namun analisis sentimen pada penelitian tersebut hanya mengelompokkan sentimen positif dan negatif saja. Berdasarkan permasalahan tersebut maka penelitian ini akan melakukan analisis sentimen hingga pada aspek-aspek layanan yang ada di hotel terutama hotel di Pekanbaru.

Berdasarkan latar belakang yang telah dijelaskan di atas maka akan dilakukan penelitian yang berjudul “Analisis Sentimen Berbasis Aspek pada Pelayanan hotel di Pekanbaru Menggunakan Metode *Naive Bayes* dan *Logistic regression*”.

1.2 Identifikasi Masalah

Adapun identifikasi masalah yang dapat diambil dari latar belakang tersebut adalah sebagai berikut.

1. Banyaknya ulasan pelanggan terkait pelayanan hotel di Pekanbaru yang beragam dan mencakup berbagai aspek (seperti pelayanan dan fasilitas) namun belum dianalisis secara optimal.
2. Kurangnya penelitian yang spesifik mengenai analisis berbasis aspek pada pelayanan hotel di Pekanbaru, terutama dalam memisahkan dan menganalisis sentimen berdasarkan aspek-aspek tertentu dari ulasan pelanggan.
3. Belum adanya perbandingan kinerja antara metode *Naïve Bayes* dan *Logistic Regression* dalam analisis sentimen berbasis aspek pada pelayanan hotel di Pekanbaru, sehingga efektivitas kedua metode tersebut dalam mengidentifikasi sentimen pada aspek-aspek spesifik ulasan pelanggan masih belum diketahui.

1.3 Rumusan Masalah

Berdasarkan masalah yang telah diidentifikasi, terdapat beberapa rumusan atau pertanyaan masalah sebagai berikut.

1. Apa saja aspek yang dibahas pada ulasan pelayanan hotel di Pekanbaru?
2. Bagaimana kinerja metode *Naïve Bayes* dan *Logistic Regression* dalam melakukan analisis sentimen berbasis aspek pada ulasan pelayanan hotel di Pekanbaru?
3. Metode mana yang paling efektif dalam menganalisis sentimen berbasis aspek pada ulasan pengunjung hotel di Pekanbaru, dilihat dari segi akurasi dan efisiensi?

1.4 Batasan Masalah

Untuk membatasi ruang lingkup penelitian, adapun batasan-batasan masalah dalam penelitian ini adalah sebagai berikut.

1. Penelitian ini hanya difokuskan pada ulasan pelanggan hotel yang berlokasi di Pekanbaru.
2. Data yang digunakan dalam penelitian ini berasal dari ulasan pelanggan yang tersedia di platform *online* yaitu Google Maps.
3. Aspek yang akan diteliti adalah aspek yang diperoleh dari pemodelan topik.
4. Penelitian ini hanya membandingkan dua metode klasifikasi teks yaitu *Naïve Bayes* dan *Logistic Regression*. Dan evaluasi kerja model dibatasi pada metrik evaluasi.

1.5 Tujuan Penelitian

Adapun tujuan yang ingin dicapai dari penelitian adalah sebagai berikut.

1. Memperoleh aspek dari pemodelan topik yang akan digunakan pada analisis sentimen berbasis aspek.
2. Membandingkan kinerja dua metode klasifikasi teks, *Naïve Bayes* dan *Logistic Regression* dalam melakukan analisis sentimen berbasis aspek pada ulasan pengunjung hotel di Pekanbaru.
3. Menentukan metode yang lebih efektif dalam hal akurasi dan efisiensi dalam menganalisis sentimen berbasis aspek pada ulasan pengunjung hotel di Pekanbaru.

1.6 Manfaat Penelitian

Manfaat yang diperoleh dari penelitian ini adalah sebagai berikut

1. Penelitian ini memberikan wawasan bagi pihak hotel di Pekanbaru dalam memahami opini pengunjung hotel secara lebih mendalam berdasarkan sentimen terhadap berbagai aspek pelayanan. Hasil analisis dapat digunakan untuk meningkatkan kualitas layanan, yang dapat meningkatkan kepuasan pengunjung.
2. Penelitian ini akan memperkaya literatur mengenai analisis sentimen berbasis aspek dalam lingkup perhotelan, khususnya di wilayah Pekanbaru.
3. Penelitian ini memberikan kontribusi dalam membandingkan efektivitas metode *Naïve Bayes* dan *Logistic Regression* dalam klasifikasi sentimen yang bisa menjadi referensi bagi penelitian di bidang yang serupa.

BAB II

LANDASAN TEORI

2.1 Tinjauan Pustaka

Studi kepustakaan dilakukan untuk memberi pengetahuan bagi penulis dalam melakukan penelitian dengan mengumpulkan data dan informasi berupa buku-buku ilmiah, skripsi, jurnal dan sumber-sumber tertulis lainnya yang telah dilakukan oleh penelitian terdahulu, sebagai berikut.

Penelitian dari Akbar et al. (2023) yang melakukan pemodelan topik pada ulasan masyarakat terhadap ulasan aplikasi PeduliLindungi. Penelitian ini menggunakan LDA dengan jumlah data yang digunakan pada penelitian ini sebanyak 13.731 data yang didapatkan dengan melakukan *scraping* pada Google Play. Hasil penelitian menunjukkan nilai koherensi yang dihitung tidak terjadi kenaikan yang signifikan setelah angka 5. Dapat disimpulkan topik utama yang dibahas pada aplikasi PeduliLindungi sebanyak 5 topik.

Penelitian oleh Gulati et al. (2022) yang melakukan komparasi 7 algoritma *machine learning* pada 72.000 *tweets*. Berdasarkan hasil yang diperoleh, algoritma Linear SVC, *Perceptron*, *Passive Aggressive Classifier*, dan *Logistic Regression* mampu mencapai akurasi maksimal di atas 98% dalam tugas klasifikasi menggunakan representasi *unigram*, *bigram*, dan *trigram*. Keempat model tersebut menunjukkan kinerja yang sangat bersaing dan hampir setara satu sama lain. Rata-rata akurasi yang dicapai oleh masing-masing model adalah 0.9815 untuk Linear SVC, 0.9765 untuk *Perceptron*, 0.9815 untuk *Passive Aggressive Classifier*, dan

0.9766 untuk *Logistic Regression*. Di sisi lain, *AdaBoost Classifier* menunjukkan performa terburuk dibandingkan model lainnya, dengan rata-rata akurasi hanya sebesar 0.7314.

Pada penelitian selanjutnya “Analisis Sentimen Berbasis Aspek pada Ulasan Pelanggan Restoran Bakso President Malang dengan Metode *Naïve Bayes Classifier*” oleh (Parasati et al., 2020). Penelitian ini menggunakan 2.152 data ulasan pelanggan dari tahun 2012 hingga tahun 2019. Data ulasan di dapatkan melalui Teknik Web Scrapping pada situs TripAdvisor dan Google Review. Analisis sentimen dalam setiap aspek menghasilkan akurasi sebesar 88% pada aspek Makanan, 76% pada aspek Layanan dan 84% pada aspek Atmosfir. Hasil klasifikasi analisis sentimen divisualisasikan dalam bentuk *dashboard* yang disertai filter berdasarkan waktu, aspek dan sentimen.

Astuti (2020) melakukan penelitian pemodelan topik analisis sentimen berbasis aspek terhadap ulasan aplikasi Tokopedia. Penelitian ini menerapkan pendekatan LDA sebagai pemodelan topik dan NB sebagai klasifikasi sentimen. Penelitian ini menggunakan data sebanyak 4.425 ulasan dalam rentang waktu 17 Oktober 2018 hingga 10 Mei 2019. Hasil dari penelitian ini menunjukkan model LDA memperoleh nilai koherensi tertinggi dengan jumlah topik sebanyak 16, namun saat divisualisasikan 16 topik tersebut saling beririsan dan berdekatan. Topik yang saling beririsan dapat dijadikan satu kluster, sehingga penelitian ini menetapkan jumlah topik sebesar 4 dengan perolehan nilai koherensi sebesar 0,53707. Sedangkan analisis sentimen berbasis aspek menghasilkan akurasi sebesar 92,5%.

Pratama et al. (2023) melakukan penelitian analisis sentimen Twitter untuk mengetahui opini publik terhadap kendaraan listrik. Data yang digunakan dalam penelitian ini adalah sebesar 1.874 data. Hasil penelitian menunjukkan *Logistic Regression* memperoleh akurasi sebesar 87,9% pada rasio *train & test set* 9:1.

Berdasarkan penelitian (Helmayanti et al., 2023) tentang Penerapan Algoritma TF-IDF dan *Naïve Bayes* untuk analisis sentimen berbasis aspek ulasan Aplikasi Flip pada Google Play Store, bahwa metode *Naïve Bayes* yang diterapkan pada ulasan Flip dengan aspek kecepatan, keamanan, dan biaya memberikan hasil terbaik dengan pembagian data 90:10 (90% data *training*, 10% data *testing*). Hasil pengujian menunjukkan akurasi rata-rata tinggi, yaitu 84% dengan rincian kecepatan 80%, keamanan 87% dan biaya 84%. Sentimen negatif paling sering muncul pada aspek kecepatan dan keamanan, sementara aspek biaya lebih didominasi sentimen positif, temuan ini memberikan wawasan bagi pengembang Flip untuk meningkatkan performa aplikasi, terutama dalam hal kecepatan dan keamanan.

Rahmawati et al. (2023) melakukan penelitian terhadap analisis sentimen penumpang terhadap pelayanan maskapai Lion Air melalui platform *online*. Penelitian ini bertujuan untuk mengetahui opini positif, negatif, dan netral dari masyarakat mengenai layanan Lion Air, mengingat tingginya volume opini yang muncul di media sosial. Tiga metode klasifikasi yang digunakan dalam penelitian ini adalah *Naïve Bayes*, *Random Forest*, dan *Logistic Regression*. Hasil menunjukkan bahwa *Logistic Regression* memberikan performa terbaik dengan akurasi 0,82. Sementara itu, metode *Naïve Bayes* dan *Random Forest* hanya menghasilkan akurasi sebesar 0,47 dan 0,39. Temuan ini menunjukkan bahwa

Logistic Regression paling efektif dalam mengklasifikasikan opini penumpang terhadap maskapai tersebut.

Tabel 2.1 Rangkuman Penelitian Terdahulu

No	Penelitian	Judul	Metode	Hasil
1	Akbar et al. (2023)	Pemodelan Topik Menggunakan Latent Dirichlet Allocation pada Ulasan Aplikasi PeduliLindungi	LDA	5 topik
2	(Parasati et al., 2020)	Analisis Sentimen Berbasis Aspek pada Ulasan Pelanggan Restoran Bakso President Malang dengan Metode <i>Naïve Bayes</i> Classifier	<i>Naïve Bayes</i>	Akurasi sebesar 88% pada aspek Makanan, 76% pada aspek Layanan dan 84% pada aspek Atmosfir.
3	(Gulati et al., 2022)	<i>Comparative analysis of machine learning-based classification models using sentiment classification of tweets related to COVID-19 pandemic</i>	Linear Svc Perceptron Passive Aggressive Classifier Logistic Regression	Akurasi $\pm 98\%$
4	(Astuti, 2020)	Analisis Sentimen Berbasis Aspek pada Aplikasi Tokopedia Menggunakan LDA dan <i>Naïve Bayes</i>	LDA & <i>Naïve Bayes</i>	4 Topik (LDA) Akurasi 92,5%

5	(Pratama et al., 2023)	Analisis Sentimen Kendaraan Listrik Pada Media Sosial Twitter Menggunakan Algoritma Logistic Regression dan Principal Component Analysis	Logistic Regression	Akurasi 87,9%
6	(Helmayanti et al., 2023)	Penerapan Algoritma TF-IDF dan <i>Naïve Bayes</i> Untuk Analisis Sentimen Berbasis Aspek Ulasan Aplikasi Flip Pada Google Play Store	<i>Naïve Bayes</i>	Akurasi aspek Kecepatan 80%, Keamanan 87% dan Biaya 84%
7	(Rahmawati et al., 2023)	Analisis Sentimen Menggunakan Algoritma Logistic Regression Pada Penerbangan Lion Air berdasarkan Ulasan Platform Online	- <i>Naïve Bayes</i> - Random Forest - Logistic Regression	- Akurasi 47% - Akurasi 39% - Akurasi 82%
8	(Nasution, A. H., & Onan, A., 2024)	Comparing the Quality of Human-Generated and LLM-Generated Annotations in Low-Resource Language NLP Tasks	Eksperimen dengan membandingkan anotasi manusia dan LLM pada tugas NLP bahasa low-resource menggunakan metrik akurasi.	notasi LLM memiliki kualitas mendekati anotasi manusia dan dapat digunakan sebagai alternatif
9	(Hanafiah, A., dkk. 2023)	Sentimen Analisis Customer Review Produk Shopee Berbasis	- NLP - Naïve Bayes - TF-IDF - wordcloud.	69% positif, 20% netral, 11% negatif; akurasi 85%.

		Wordcloud dengan Algoritma Naïve Bayes Classifier		
10	(Nasution, A. H., dkk. 2023).	Benchmarking Open-Source Large Language Models for Sentiment and Emotion Classification in Indonesian Tweets	Benchmarking beberapa LLM open-source pada tugas klasifikasi sentimen dan emosi.	Model LLM open-source mencapai lebih dari 93–94% performa ChatGPT-4 pada klasifikasi sentimen dan emosi bahasa Indonesia dalam skenario zero-shot.

Berdasarkan tinjauan pustaka, analisis sentimen dapat dilakukan menggunakan berbagai pendekatan, termasuk *Large Language Models* (LLM) yang memiliki performa tinggi, namun masih memiliki keterbatasan dari sisi kompleksitas, kebutuhan sumber daya komputasi, serta potensi *data leakage*. Oleh karena itu, penelitian **“Analisis Sentimen Berbasis Aspek pada Pelayanan Hotel di Pekanbaru Menggunakan Metode Naïve Bayes dan Logistic Regression”** memilih metode *Naïve Bayes* dan *Logistic Regression* karena bersifat sederhana, efisien, mudah diimplementasikan, serta sesuai untuk data teks berbahasa Indonesia. Penggunaan dua metode ini juga memungkinkan perbandingan performa guna menentukan metode yang paling optimal dalam menganalisis sentimen berbasis aspek pada ulasan pelayanan hotel.

2.2 Dasar Teori

2.2.1 *Text Mining*

Text mining merupakan proses menemukan dan mengekstrak informasi dari data tak terstruktur atau tersirat dalam dokumen teks menjadi lebih eksplisit (Glanville, 2018). Dua metode dari penerapan *text mining* dalam menemukan dan mengekstrak informasi adalah pemodelan topik dan analisis sentimen (Kochmar, 2022; Liu, 2020; Teng & Zhang, 2021).

2.2.1.1 Analisis Sentimen

Analisis sentimen atau *opinion mining* merupakan disiplin ilmu dalam pemrosesan bahasa alamiah yang berfokus mempelajari metode untuk menganalisis opini, sentimen, atau emosi seseorang terhadap suatu entitas (Liu, 2020). Analisis sentimen merujuk pada proses mengklasifikasikan opini, sentimen, atau emosi ke dalam kategori positif, negatif, atau netral dari data teks. Tujuan dari analisis sentimen adalah untuk memahami pandangan orang terhadap berbagai hal, seperti produk, layanan, atau topik tertentu di media sosial (Lamba & Madhusudhan, 2022).

Sementara itu, analisis sentimen berbasis aspek merupakan sub area dari analisis sentimen yang meninjau opini dan sasarannya (aspek) (Liu, 2020). Fokus dari analisis sentimen berbasis aspek adalah mengklasifikasikan sentimen terhadap aspek tertentu. Dengan itu, analisis sentimen berbasis aspek dapat membantu dalam membuat ringkasan opini tentang suatu hal atau entitas dengan mempertimbangkan secara spesifik opini tersebut dan aspeknya.

2.2.1.2 Pemodelan Topik

Pemodelan topik adalah salah satu pendekatan *text mining* yang digunakan untuk menemukan topik dari dokumen teks dengan skala yang besar (Liu, 2020). Topik merupakan gagasan utama atau pokok pembicaraan yang dibahas dalam sebuah teks. Menurut Lamba & Madhusudhan (2022), pemodelan topik dapat mengelompokkan dokumen berdasarkan pola, frekuensi, dan jarak serta dapat membantu dalam proses anotasi data. Topik dalam konteks pemrosesan bahasa alamiah juga dapat disebut sebagai aspek (Liu, 2020), sehingga proses anotasi aspek pada dokumen dapat dilakukan berdasarkan topik-topik yang telah diidentifikasi melalui pemodelan topik.

2.2.1.3 Text Generation

Text generation (generasi teks) adalah proses pembuatan teks secara otomatis yang menyerupai cara manusia menulis atau berbicara. Ini merupakan salah satu topik utama dalam *Natural Language Processing* (NLP), yaitu cabang kecerdasan buatan (AI) yang fokus pada pemrosesan dan pemahaman bahasa manusia (J. Li et al., 2022).

Pada dasarnya, tujuan dari *text generation* adalah untuk menghasilkan teks yang koheren dan terbaca, yang dapat digunakan untuk berbagai keperluan. Input yang digunakan untuk menghasilkan teks ini bisa berupa berbagai jenis data, seperti teks, gambar, dan tabel.

Teknik *text generation* telah mengalami perkembangan pesat dalam beberapa dekade terakhir, terutama dengan kemunculan model bahasa berbasis *deep learning* seperti GPT (*Generative Pre-trained Transformer*) yang dapat menghasilkan teks

yang sangat mirip dengan tulisan manusia. Model ini dapat digunakan untuk berbagai aplikasi, mulai dari *chatbots* yang dapat menjawab pertanyaan pengguna, penulisan otomatis untuk artikel atau laporan, hingga penerjemahan otomatis.

2.2.2 Ulasan

Ulasan merupakan sebuah evaluasi atau pandangan yang diberikan individu atau pengguna suatu produk, layanan, tempat, atau topik tertentu. Ulasan dapat diberikan untuk berbagai tujuan, seperti memberikan umpan balik kepada penyedia layanan atau produk, memberikan panduan kepada pengguna lain dalam proses pengambilan keputusan pembelian atau hanya berbagi pengalaman. Secara khusus, ulasan *online* memberikan informasi tentang penilaian yang biasanya dinyatakan dalam bentuk *rating*, sementara pengalaman diungkapkan melalui komentar atau saran (Baskoro, 2021)

2.2.3 Google Maps

Google Maps (maps.google.com) adalah platform pemetaan web yang diluncurkan oleh Google pada tahun 2005. Google Maps menampilkan citra satelit, foto udara, peta jalan, pemandangan panorama jalan 360° yang interaktif (Street View), kondisi lalu lintas waktu terkini, dan perencanaan rute untuk bepergian dengan berjalan kaki, mobil, dan transportasi umum. Pada tahun 2020, Google Maps digunakan oleh lebih dari satu miliar orang setiap bulan di seluruh dunia¹².

Google Maps dapat digolongkan sebagai media sosial situs ulasan produk/layanan karena pada platform tersebut pengguna dapat berbagi ulasan mengenai layanan suatu tempat yang terdapat di Google Maps. Sementara itu,

sebagai situs jejaring sosial, pengguna pada platform Google Maps dapat membuat profil, mengikuti pengguna lain, dan berbagi foto atau video di suatu tempat yang terdapat di Google Maps, mereka juga dapat mengajukan dan menjawab pertanyaan tentang layanan atau produk dari suatu tempat.

2.2.4 Web Scrapping

Web scraping adalah cara yang digunakan untuk mendapatkan data atau informasi dari suatu website yang dilakukan secara otomatis. Tujuan dari web scraping adalah untuk menggali informasi dari situs web yang berbeda dan tidak terstruktur kemudian mengubahnya menjadi bentuk yang lebih rapi dan terstruktur dalam bentuk basis data, spreadsheets, atau Comma Separated Values (CSV). sebuah teknik untuk mengekstrak data dari sebuah website atau sistem tertentu. Data scraping biasanya juga disebut dengan data extraction. data extraction adalah sebuah hal yang digunakan untuk beberapa pekerjaan yang terkait dengan digital marketing, misalnya riset konten. Salah satu cara yang dapat digunakan untuk melakukan ekstraksi data adalah dengan memanfaatkan Application Programming Interface (API). API memungkinkan untuk mengakses sebuah situs dengan format data yang lebih terstruktur.

Pada tahap pengumpulan data, penulis mengumpulkan data ulasan di Rumah sakit swasta di Pekanbaru menggunakan script Python untuk scrape data teks di halaman web yang dijalankan di Jupyter Notebook.

2.2.5 Pra-pemrosesan Data

Pra-pemrosesan merupakan proses transformasi teks sebelum dianalisis dengan tujuan untuk mengurangi jumlah kosakata yang tidak relevan dan menghilangkan *noise* (gangguan) (Hickman et al., 2022). Proses pra-pemrosesan data meliputi *case folding*, koreksi ejaan dan singkatan, penghapusan emoji, angka, tanda baca, *whitespace*, karakter tunggal, kata tunggal, dan kalimat berbahasa Inggris, tokenisasi, penghapusan *stopwords*, dan *stemming* dijelaskan sebagai berikut (Anandarajan et al., 2019b).

1. *case folding* (mengubah semua huruf menjadi huruf kecil), koreksi ejaan dan singkatan, penghapusan emoji, angka, tanda baca, *whitespace*, karakter tunggal, kata tunggal, dan kalimat berbahasa Inggris merupakan tahapan yang berfungsi membersihkan data dari elemen yang tidak relevan atau mengganggu dan membuat data menjadi konsisten dan seragam.
2. Tokenisasi, proses mengubah teks menjadi unit-unit yang lebih kecil, seperti kata, kelompok kata, atau frasa.
3. Penghapusan Stopwords, merupakan proses penghapusan kata-kata penghubung yang tidak memiliki arti atau makna khusus dalam sebuah teks.
4. Stemming, merupakan proses mengurai kata-kata dalam teks menjadi bentuk dasar dari kata tersebut.

2.2.6 Teknik *Prompting*

Menurut Fagbohun et al. (2023), *prompt* adalah jembatan penerjemah antara komunikasi manusia dalam bahasa alami dan kemampuan komputasional model bahasa besar (LLM) seperti ChatGPT. Dalam praktiknya, *prompt* terdiri dari

instruksi berbahasa alami yang berfungsi sebagai bahasa perantara, menerjemahkan permintaan manusia menjadi tugas-tugas yang dapat dieksekusi oleh mesin. Dua di antara berbagai teknik *prompting* yang ada adalah *zero-shot* dan *few-shot*.

2.2.6.1 Zero-Shot

Menurut Schulhoff et al. (2025), *zero-shot prompting* adalah teknik di mana model diberikan instruksi langsung dalam bentuk bahasa alami tanpa contoh apa pun. Artinya, *prompt* yang digunakan tidak menyertakan demonstrasi atau contoh tugas sebelumnya. Model hanya mengandalkan pemahamannya yang sudah dimiliki dari proses pelatihan sebelumnya untuk menafsirkan dan menjalankan instruksi tersebut.

2.2.6.2 Few-Shot

Menurut Schulhoff et al. (2025), *Few-shot prompting* adalah teknik dalam interaksi dengan model bahasa besar (LLM) di mana pengguna memberikan beberapa contoh dalam *prompt* untuk menunjukkan cara menyelesaikan suatu tugas. Contoh-contoh ini digunakan sebagai demonstrasi, sehingga model dapat belajar secara langsung dari konteks yang diberikan dan menerapkannya saat menjawab permintaan baru yang serupa.

Teknik ini memanfaatkan kemampuan *in-context learning*, yaitu kemampuan model untuk mengenali pola dari contoh-contoh dalam *prompt* tanpa perlu pelatihan ulang (*training*). *Few-shot prompting* sering digunakan untuk meningkatkan akurasi model pada tugas-tugas yang kompleks atau ambigu, yang sulit dijawab dengan instruksi saja (*zero-shot*).

2.2.7 Imbalanced Data

Imbalance data terjadi ketika jumlah sampel data untuk suatu kelas (kelas mayoritas) jauh lebih banyak dibandingkan dengan kelas lainnya (kelas minoritas). Imbalance data dapat menyebabkan model pembelajaran mesin tidak akurat dalam memprediksi kelas minoritas serta dapat mengurangi akurasi model (Wongvorachan et al., 2023). Salah satu cara dalam menangani ketidakseimbangan data adalah dengan cara menambah jumlah sampel dari kelas minoritas atau dapat disebut *over-sampling* (Zhihao et al., 2019). *Over-sampling* memiliki keuntungan dalam penyeimbangan data dengan tidak mengurangi informasi yang mungkin berguna dalam meningkatkan akurasi model klasifikasi (Wongvorachan et al., 2023; Zhihao et al., 2019). Salah satu metode *over-sampling* adalah *Synthetic Minority Over-sampling Technique* (SMOTE). *Synthetic Minority Over-sampling Technique* merupakan metode yang bertujuan untuk menyeimbangkan jumlah sampel antara kelas minoritas dan mayoritas dengan cara membuat sampel sintetis baru untuk kelas minoritas (Chawla et al., 2002).

2.2.8 Ekstraksi Fitur

Ekstraksi fitur merupakan proses yang merepresentasikan pesan teks dengan mengidentifikasi dan memilih fitur teks yang relevan guna mengurangi dimensi ruang fitur dengan menghapus fitur yang tidak berkorelasi atau berlebihan (Liang et al., 2017). Ekstraksi fitur juga merujuk pada representasi fitur dalam bentuk numerik agar dapat dipahami oleh mesin, sehingga mempengaruhi akurasi pada model klasifikasi (Dzisevic & Sesok, 2019; Lamba & Madhusudhan, 2022).

Berbagai metode dapat digunakan dalam membangun ekstraksi fitur, dua di antaranya dijelaskan sebagai berikut.

2.2.8.1 N-Gram

N-gram merupakan urutan unit kata atau karakter dengan panjang tertentu yang dilambangkan dengan n (Kochmar, 2022). Menurut Anandarajan et al. (2019b), n-gram adalah cara lain untuk memecah teks menjadi bagian-bagian yang lebih kecil dalam proses tokenisasi. Misalnya, *unigram* merupakan urutan satu kata ($n = 1$), *bigram* merupakan urutan dua kata ($n = 2$), dan *trigram* merupakan urutan tiga kata ($n = 3$), dan seterusnya (Anandarajan et al., 2019b; Kochmar, 2022). Contoh penerapan n-gram, seperti “*saya suka minum teh*” menjadi “*saya*”, “*suka*”, “*minum*”, “*teh*” untuk $n = 1$, “*saya suka*”, “*suka minum*”, “*minum teh*” untuk $n = 2$.

2.2.8.2 TF-IDF

TF-IDF atau *Term Frequency-Inverse Document Frequency* merupakan metode pembobotan kata-kata dalam sebuah dokumen berdasarkan seberapa sering kata tersebut muncul dalam dokumen tersebut, dan seberapa penting kata tersebut di dalam dokumen (Anandarajan et al., 2019a; Dzisevic & Sesok, 2019). Proses perhitungan menggunakan TF-IDF ditunjukkan pada persamaan (2.1).

$$TF - IDF = TF * IDF \quad (2,1)$$

TF atau *Term Frequency* merupakan frekuensi kemunculan kata (t) dalam sebuah dokumen (d) (Lamba & Madhusudhan, 2022). TF direpresentasikan pada persamaan (2.2).

$$TF(t, d) = \frac{\text{jumlah kata } k \text{ muncul di dokumen}}{\text{jumlah semua kata pada dokumen}} \quad (2.2)$$

IDF atau *Inverse Document Frequency* merupakan pengukuran signifikansi sebuah kata dalam suatu korpus, dengan memprioritaskan kata yang jarang muncul karena cenderung lebih informatif dan memberikan nilai yang lebih besar dalam proses analisis teks (Lamba & Madhusudhan, 2022). IDF direpresentasikan pada persamaan (2.3).

$$IDF = \log \left(\frac{\text{jumlah dokumen dalam korpus}}{\text{jumlah dokumen yang terdapat kata } k} \right) \quad (2.3)$$

2.2.9 Algoritma Text Mining

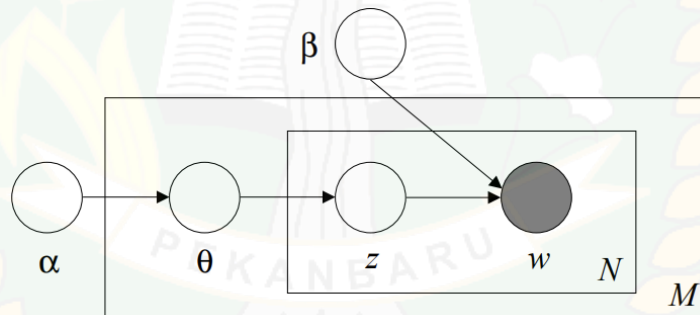
2.2.9.1 Latent Dirichlet Allocation

Latent Dirichlet Allocation (LDA) adalah sebuah model pembelajaran tanpa pengawasan (*unsupervised learning*) yang menerapkan pendekatan probabilistik untuk menganalisis dokumen. Model ini berasumsi bahwa setiap dokumen tersusun dari campuran beberapa topik, di mana setiap topik merepresentasikan distribusi probabilitas atas sejumlah kata (Liu, 2020). LDA merupakan jaringan Bayesian yang digunakan untuk melakukan pemodelan topik terhadap data teks (Teng & Zhang, 2021).

LDA pertama kali diperkenalkan oleh Blei et al. (2003) sebagai model probabilitas generatif untuk secara otomatis menemukan topik tersembunyi dari data teks. Perhitungan probabilitas LDA adalah sebagai berikut.

$$p(D|\alpha, \beta) = \prod_{d=1}^M \int p(\theta_d|\alpha) \left(\prod_{n=1}^{N_d} \sum_{z_{dn}} p(z_{dn}|\theta_d) p(w_{dn}|z_{dn}, \beta) \right) d\theta_d \quad (2.4)$$

Pada persamaan (2.1), untuk memperoleh kemungkinan pemodelan, baik α maupun β merepresentasikan hyperparameter. Parameter Dirichlet prior untuk distribusi subjek di tingkat dokumen diwakili oleh simbol α , sedangkan distribusi probabilitas kata untuk subjek tertentu diwakili oleh simbol β . Distribusi topik dokumen d adalah vektor θ . Dokumen d memiliki topik tersembunyi, yang ditunjukkan oleh notasi z . Notasi M dan N menunjukkan panjang dokumen dan jumlah istilah yang ada di dalamnya. Notasi grafis dari model LDA direpresentasikan pada Gambar 2.1.



Gambar 2.1 Notasi Grafis Model LDA
(Sumber: Blei et al. (2003))

2.2.9.2 *Naïve Bayes*

Naïve Bayes merupakan metode klasifikasi yang menggunakan probabilitas dan peluang dari prinsip Bayesian untuk memprediksi peluang di masa depan berdasarkan pengalaman di masa sebelumnya (Yang, 2018). Menurut Ray (2019) *Naïve Bayes* dapat dengan mudah diimplementasikan dan memberikan kinerja yang

baik, terutama dengan data latih yang lebih kecil serta memiliki kemampuan dalam menangani data diskrit dan kontinu. Persamaan prinsip Bayesian memiliki bentuk umum sebagai berikut.

$$P(H|X) = \frac{P(X|H) \cdot P(H)}{P(X)} \quad (2.5)$$

Keterangan:

- X : Bukti atau fakta
- H : Hipotesis
- $P(H|X)$: Probabilitas *posterior* H dengan syarat X
- $P(X|H)$: Probabilitas *posterior* X dengan syarat H
- $P(H)$: Probabilitas prior hipotesis H
- $P(X)$: Probabilitas prior bukti X

2.2.9.3 Logistic Regression

Logistic Regression merupakan metode klasifikasi yang menentukan variabel dependen diskrit dari himpunan variabel independen, serta bertujuan untuk memprediksi probabilitas terjadinya suatu peristiwa dengan menyesuaikan data ke dalam fungsi logit dari kurva logistik (Arianto & Budi, 2023). *Logistic Regression* memiliki keunggulan dalam kemudahan implementasi dan kecepatan pemrosesan (Ray, 2019). Bentuk umum persamaan regresi logistik dapat dilihat sebagai berikut.

$$\pi(x) = \frac{\exp(\alpha + \beta x)}{1 + \exp(\alpha + \beta x)} \quad (2.6)$$

Keterangan:

$\pi(x)$: Probabilitas prediksi kejadian

α : Intersep (konstanta)

β : Koefisien variabel independen (x)

2.2.10 *K-Fold Cross Validation*

Cross validation merupakan teknik validasi model pembelajaran mesin yang bertujuan untuk mengestimasi tingkat kesalahan uji dari suatu model dengan membagi data menjadi *subset* latih dan uji yang berbeda (James et al., 2021; Lamba & Madhusudhan, 2022). Salah satu metode *cross validation* adalah *k-fold cross validation*.

K-fold cross validation merupakan teknik evaluasi model pembelajaran mesin yang melibatkan pembagian acak dari data latih dan uji menjadi k bagian dengan ukuran yang sama (James et al., 2021). *K-fold cross validation* memberikan gambaran yang lebih komprehensif tentang seberapa baik model akan berkinerja pada data baru yang independen (Lamba & Madhusudhan, 2022). Metode ini mengulangi proses sebanyak k kali, pada setiap perulangan data dibagi menjadi k bagian dengan satu bagian digunakan untuk validasi atau pengujian dan $k - 1$ bagian lainnya digabungkan menjadi *subset* pelatihan untuk evaluasi model. Alur kerja *k-fold cross validation* dengan $k = 5$ diilustrasikan pada Gambar 2.2.

Perulangan ke-1	Bagian 1 (Uji)	Bagian 1	Bagian 1	Bagian 1	Bagian 1
Perulangan ke-2	Bagian 2	Bagian 2 (Uji)	Bagian 2	Bagian 2	Bagian 2
Perulangan ke-3	Bagian 3	Bagian 3	Bagian 3 (Uji)	Bagian 3	Bagian 3
Perulangan ke-4	Bagian 4	Bagian 4	Bagian 4	Bagian 4 (Uji)	Bagian 4
Perulangan ke-5	Bagian 5	Bagian 5	Bagian 5	Bagian 5	Bagian 5 (Uji)

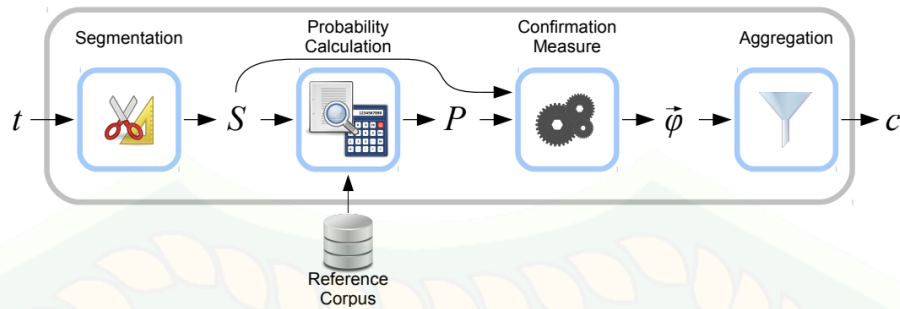
Gambar 2.2 Alur Kerja *K-Fold Cross Validation*

2.2.11 Metrik Evaluasi

Tahap evaluasi merupakan faktor penting dalam menilai kinerja model. Tahap evaluasi memberikan gambaran objektif tentang sejauh mana model dapat diandalkan. Adapun dua tahap evaluasi mengenai pemodelan topik dan model klasifikasi analisis sentimen berbasis aspek dijabarkan sebagai berikut.

2.2.11.1 *Topic Coherence*

Dalam pemodelan topik, *Latent Dirichlet Allocation* akan dievaluasi berdasarkan *topic coherence*, yaitu suatu ukuran yang mengukur seberapa informatif topik yang dihasilkan oleh model dengan mengamati kumpulan kata dalam topik tersebut (Kherwa & Bansal, 2020; Korencic et al., 2021; Tijare & Rani, 2020). Metrik *topic coherence* yang digunakan dalam penelitian ini adalah metrik “c_v” yaitu *coherence score* yang diperoleh dari pengelompokan sekelompok kata-kata teratas dan evaluasi tidak langsung menggunakan metode *normalized pointwise mutual information* dan *cosine similarity* (Kapadia, 2019).



Gambar 2.3 Proses Perhitungan *Coherence Score*
(Sumber: (Röder et al., 2015))

Gambar 2.3 merupakan gambaran dari proses perhitungan *coherence score*. Pertama, set kata t dari model topik dipecah menjadi subset pasangan kata S . Kemudian, probabilitas kata P dihitung menggunakan korpus referensi yang diberikan. Kedua, baik subset pasangan kata S maupun probabilitas P yang telah dihitung digunakan oleh *Confirmation Measure* untuk menghitung vektor kesepakatan ϕ dari pasangan S . Akhirnya, nilai-nilai ini digabungkan untuk membentuk nilai koherensi tunggal c .

2.2.11.2 *Perplexity*

Perplexity adalah ukuran dalam teori informasi yang digunakan untuk mengevaluasi kesamaan antara distribusi probabilitas (m) dengan distribusi asli (p) (Arora & Rangarajan, 2016). Nilai *perplexity* dapat dihitung sebagai kebalikan dari rata-rata probabilitas pada set pengujian T di mana n adalah jumlah kata dalam set pengujian T .

$$Perplexity(T) = \frac{1}{\sqrt[n]{m(T)}} \quad (2.7)$$

2.2.12 *Elbow Method*

Elbow method adalah teknik yang digunakan untuk mengidentifikasi titik tertentu pada grafik di mana terjadi penurunan signifikan, menghasilkan lengkungan tajam yang menyerupai bentuk siku (Syahfitri et al., 2023). Dengan begitu, topik yang dihasilkan menjadi lebih relevan, terstruktur, dan informatif untuk dianalisis lebih lanjut.

Proses ini dilakukan menggunakan teknik yang mengevaluasi perubahan nilai *coherence score* dan *perplexity* saat jumlah topik ditambah. *Coherence score* digunakan untuk mengukur sejauh mana topik yang dihasilkan dapat dipahami atau diinterpretasikan dengan baik, sedangkan *perplexity* mengukur seberapa baik model mampu memprediksi data yang digunakan.

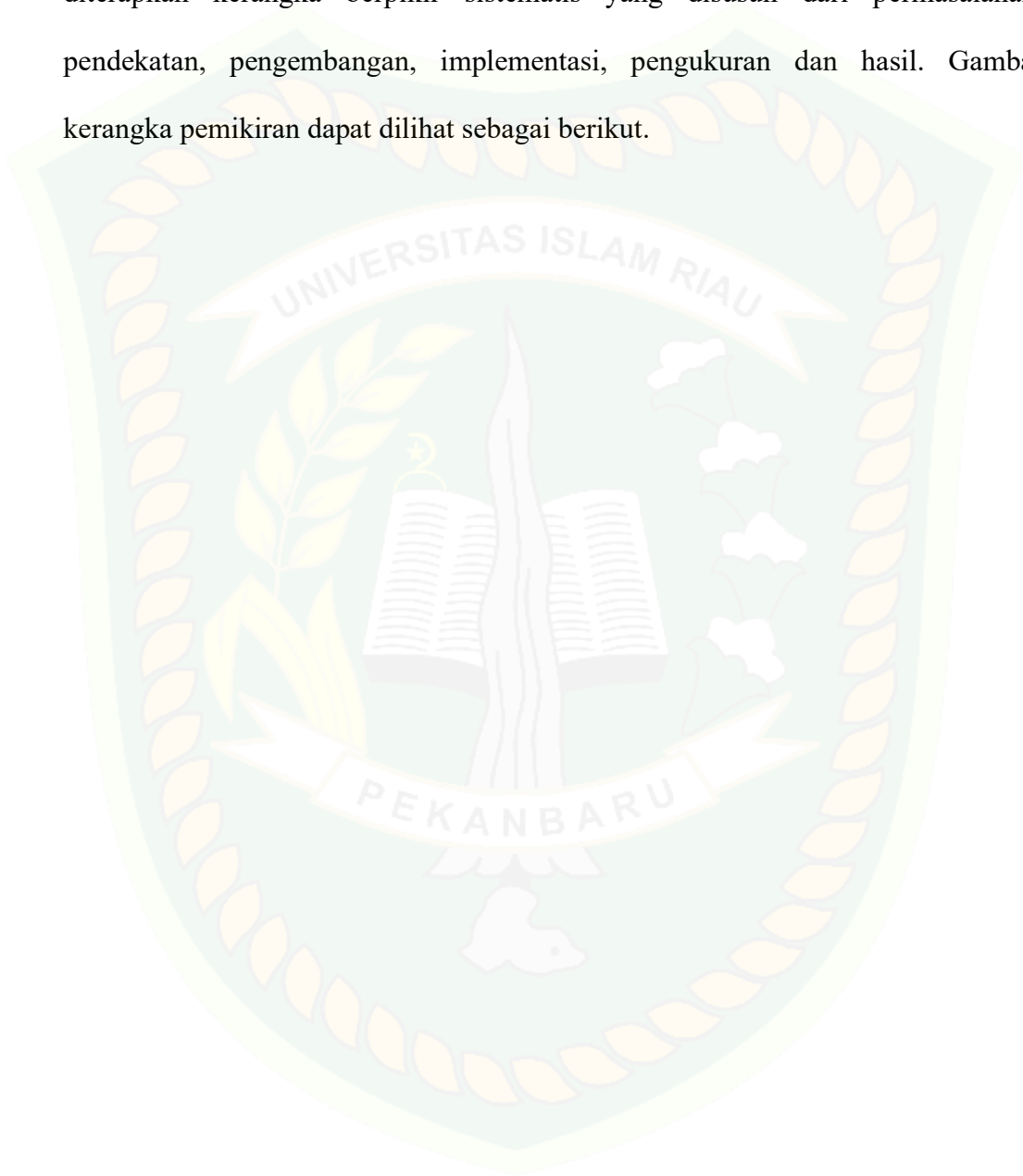
2.2.12.1 Akurasi

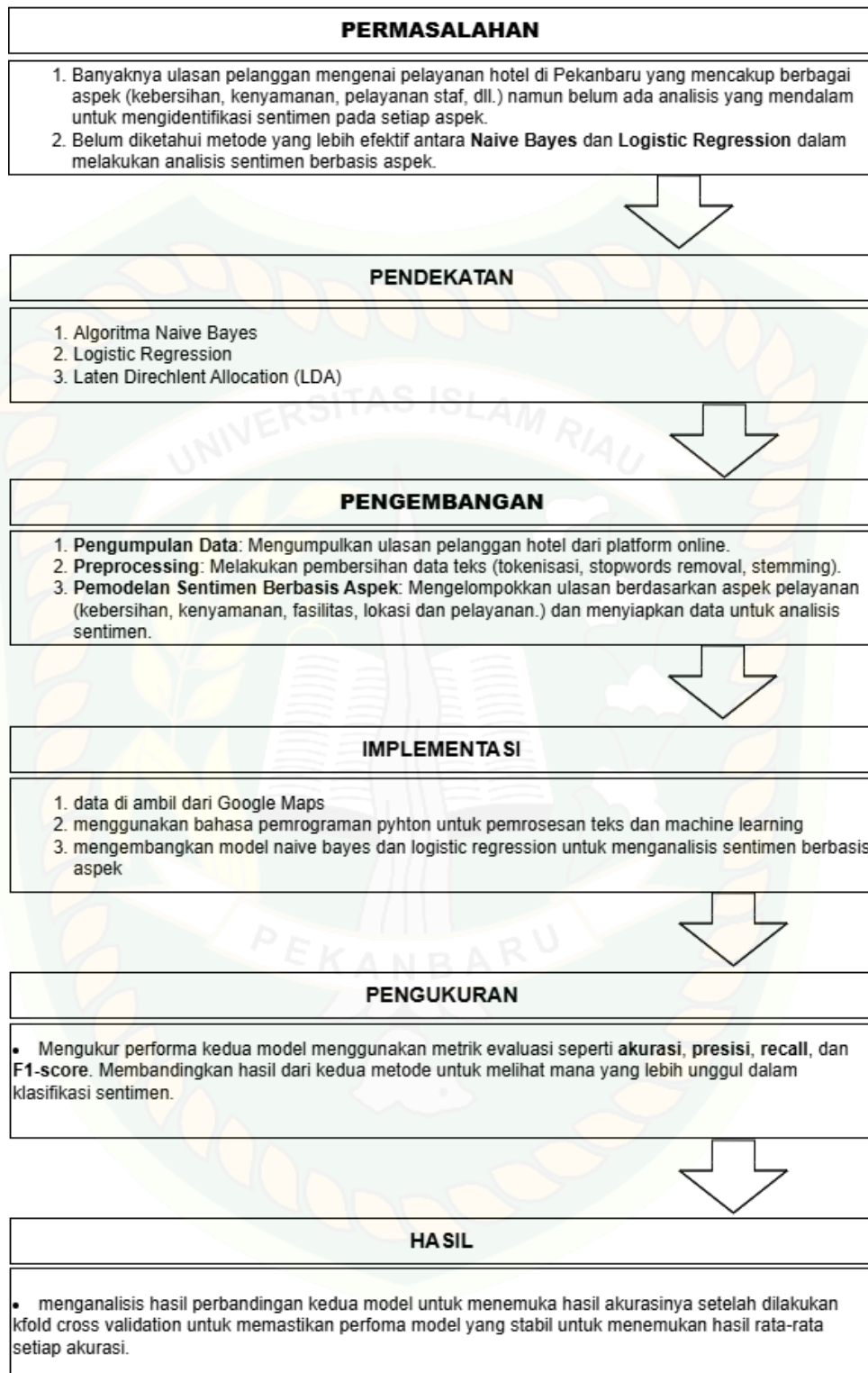
Dalam menilai kinerja model klasifikasi, terdapat beberapa metode yang dapat dijadikan metrik evaluasi salah satunya adalah Akurasi. Akurasi merupakan ukuran evaluasi yang mengindikasikan seberapa tepat model dalam membuat prediksi yang benar dari seluruh prediksi yang dilakukan (Narkhede, 2018). Akurasi dihitung berdasarkan jumlah dari *True Positive* dan *True Negative* dan dibagi dengan total data (Vujović, 2021).

$$Akurasi = \frac{(TP + TN)}{(TP + FP + TN + FN)} \quad (2.8)$$

2.3 Kerangka Pemikiran

Untuk memecahkan masalah yang sedang diteliti dalam penelitian ini, diterapkan kerangka berpikir sistematis yang disusun dari permasalahan, pendekatan, pengembangan, implementasi, pengukuran dan hasil. Gambar kerangka pemikiran dapat dilihat sebagai berikut.





Gambar 2.4 Kerangka Pemikiran

BAB III

METODOLOGI PENELITIAN

3.1 Tempat Penelitian

Dalam penelitian ini, penulis mengambil data dari platform *online* yaitu Google Maps, data yang diambil berupa data ulasan hotel yang ada di Pekanbaru. Penarikan data diambil dengan cara web *scrapping* menggunakan perpustakaan Python yaitu *Beautifulsoup*. Informasi yang terkumpul dari ulasan tersebut sangat bermanfaat untuk memahami latar belakang data yang akan penulis analisis lebih lanjut.

3.2 Alat dan Bahan

Alat dan bahan yang digunakan dalam penelitian ini adalah sebagai berikut.

1. Perangkat Keras (*Hardware*)

Perangkat keras yang digunakan dalam penelitian ini adalah satu unit laptop berikut spesifikasi teknisnya:

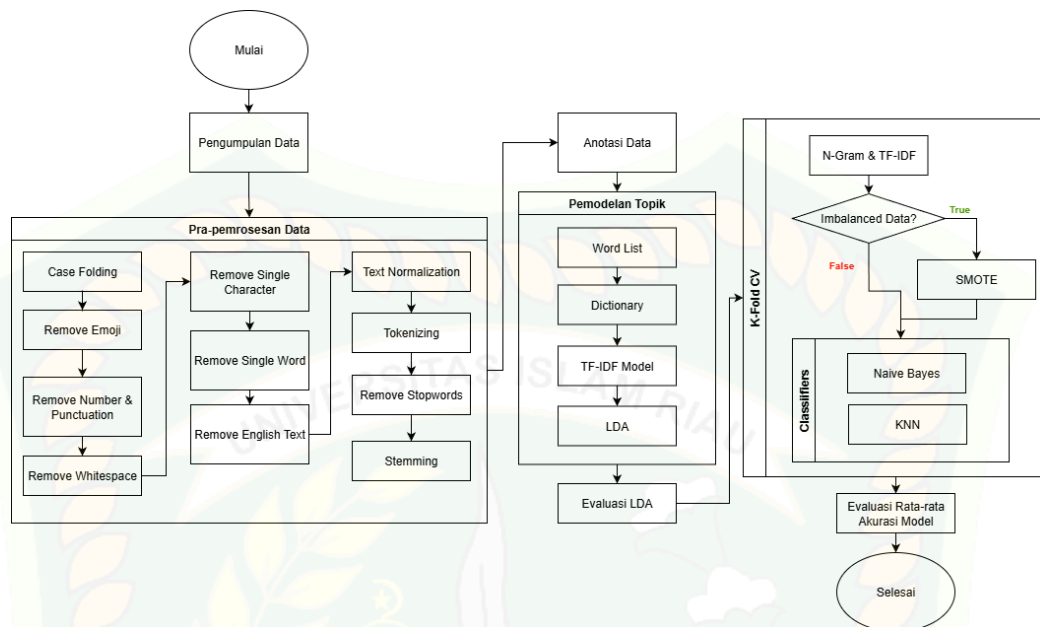
- a. *Processor* Intel(R) Core™ i5-8250U dengan kecepatan 1.8 GHz
- b. Memori instalasi sebesar 8 GB
- c. Sistem operasi 64-bit *Windows* 10

2. Perangkat lunak (*Software*)

Perangkat lunak yang terlibat dalam penelitian ini meliputi:

- a. *Jupyter Notebook* untuk membuat dan mengedit sel kode
- b. Bahasa pemrograman *Python*
- c. *Microsoft Edge* sebagai penjelajah web

3.3 Tahapan Penelitian



Gambar 3.1 Gambaran Umum Penelitian

Tahapan penelitian untuk “Analisis Sentimen berbasis Aspek pada Pelayanan Hotel di Pekanbaru menggunakan Metode *Naïve bayes* dan *Logistic Regression*” dimulai dengan pengumpulan data dari *Google Maps*. Data yang diambil berupa ulasan pelanggan hotel yang ada di Pekanbaru. Selanjutnya dilakukan tahapan pra-pemrosesan data, yang mencakup *case folding*, penghapusan angka dan tanda baca, penghapusan spasi berlebih, penghapusan karakter tunggal, penghapusan kata tunggal, penghapusan kalimat berbahasa Inggris, normalisasi teks, tokenisasi, penghapusan *stopwords*, dan *stemming*.

Tahapan selanjutnya adalah melakukan anotasi data ke dalam kelas positif atau negatif. Tahapan anotasi data dilakukan dengan memanfaatkan *text generation large language models* yaitu GPT. Berikutnya adalah melakukan pemodelan topik menggunakan *Latent Dirichlet Allocation* yang akan dievaluasi dan dianalisis

menggunakan *elbow method* berdasarkan *coherence score* dan *perplexity* guna memperoleh aspek yang akan dilanjutkan ke tahap analisis sentimen berbasis aspek. Pelatihan dan pengujian model akan dilakukan menggunakan K-Fold dengan K=10. Kemudian dilanjutkan ke tahap ekstraksi fitur, yaitu N-Grams dan TF-IDF. Sebelumnya, jika polaritas sentimen pada tiap aspek seimbang, maka data akan dilakukan SMOTE. Tahapan analisis sentimen berbasis aspek akan dievaluasi berdasarkan akurasi model.

3.3.1 Pengumpulan Data

Pengumpulan data adalah tahap awal dalam melakukan analisis sentimen dan pemodelan topik. Data yang digunakan penulis dalam penelitian ini adalah data ulasan dari pengguna Google Maps di hotel yang ada di Pekanbaru. Penulis mengumpulkan data ulasan Google Maps menggunakan teknik scraping dengan memanfaatkan pustaka BeautifulSoup.

3.3.2 Pra-pemrosesan Data

Sebelum melakukan pemodelan topik, tahapan pra-pemrosesan data yang akan dilakukan pada penelitian ini meliputi *case folding*, penghapusan emoji, angka, tanda baca, *whitespace*, dan karakter tunggal, penghapusan kata tunggal, penghapusan kalimat berbahasa Inggris, tokenisasi, penghapusan *stopwords*, dan *stemming*.

3.3.2.1 Case Folding

Tahap pra-pemrosesan pertama yang dilakukan pada penelitian ini adalah melakukan *case folding* dengan tujuan membuat kata-kata dalam teks menjadi

huruf kecil. Proses ini dilakukan karena berguna untuk pembobotan kata pada ekstraksi fitur. Perbandingan sebelum dan sesudah *case folding* dapat dilihat pada Tabel 3.1.

Tabel 3.1 *Case Folding*

Sebelum	Sesudah
Lokasi pas,,, pelayanan ramah ✨	lokasi pas,,, pelayanan ramah ✨
Roomnya sgt nyaman dan bersih	roomnya sgt nyaman dan bersih
1. Lokasi hotel strategis 2. pelayanan hotel sangat memuaskan..	1. lokasi hotel strategis 2. pelayanan hotel sangat memuaskan..

3.3.2.2 Penghapusan Emoji, Angka, dan Tanda Baca

Tahap pra-pemrosesan selanjutnya adalah penghapusan emoji, angka, dan tanda baca. Perbandingan sebelum dan sesudah pada tahapan ini dapat dilihat pada Tabel 3.2.

Tabel 3.2 Penghapusan Emoji, Angka, dan Tanda Baca

Sebelum	Sesudah
lokasi pas,,, pelayanan ramah ✨	lokasi pas pelayanan ramah
roomnya sgt nyaman dan bersih	roomnya sgt nyaman dan bersih
1. lokasi hotel strategis 2. pelayanan hotel sangat memuaskan..	lokasi hotel strategis pelayanan hotel sangat memuaskan

3.3.2.3 Penghapusan *Whitespace* dan Karakter Tunggal

Tahap selanjutnya adalah penghapusan *whitespace* dan karakter tunggal. Perbandingan sebelum dan sesudah pada tahapan ini dapat dilihat pada Tabel 3.3.

Tabel 3.3 Penghapusan *Whitespace* dan Karakter Tunggal

Sebelum	Sesudah
lokasi pas pelayanan ramah	lokasi pas pelayanan ramah
roomnya sgt nyaman dan bersih	roomnya sgt nyaman dan bersih

lokasi hotel strategis pelayanan hotel sangat memuaskan	lokasi hotel strategis pelayanan hotel sangat memuaskan
---	---

3.3.2.4 Penghapusan Kata Tunggal

Tahap ini digunakan untuk menyaring dan menghapus teks atau dokumen yang memiliki jumlah kata kurang dari dua. Tujuan utama dari tahap ini adalah untuk menghilangkan teks yang terlalu pendek dan kemungkinan tidak mengandung informasi yang cukup relevan untuk analisis lebih lanjut.

3.3.2.5 Penghapusan Kalimat Berbahasa Inggris

Tahap ini berfungsi menyaring atau menghapus kalimat yang sebagian besar atau seluruhnya menggunakan bahasa Inggris, sehingga hanya ulasan berbahasa Indonesia yang dipertahankan untuk keperluan analisis lebih lanjut.

3.3.2.6 Koreksi Ejaan dan Singkatan

Tahap ini merupakan tahap normalisasi ejaan dan singkatan menjadi bahasa formal. Tahap ini menggunakan kamus dari Suciati & Budi (2019) (Arianto & Budi, 2023). Perbandingan sebelum dan sesudah koreksi ejaan dan singkatan dapat dilihat pada Tabel 3.4.

Tabel 3.4 Koreksi Ejaan dan Singkatan

Sebelum	Sesudah
lokasi pas pelayanan ramah	lokasi pas pelayanan ramah
roomnya sgt nyaman dan bersih	roomnya sangat nyaman dan bersih
lokasi hotel strategis pelayanan hotel sangat memuaskan	lokasi hotel strategis pelayanan hotel sangat memuaskan

3.3.2.7 Tokenisasi

Tahap selanjutnya adalah tokenisasi, yaitu proses memecah kalimat menjadi unit yang lebih kecil. Perbandingan sebelum dan sesudah tokenisasi dapat dilihat pada Tabel 3.5.

Tabel 3.5 Tokenisasi

Sebelum	Sesudah
lokasi pas pelayanan ramah	['lokasi', 'pas', 'pelayanan', 'ramah']
roomnya sangat nyaman dan bersih	['roomnya', 'sangat', 'nyaman', 'dan', 'bersih']
lokasi hotel strategis pelayanan hotel sangat memuaskan	['lokasi', 'hotel', 'strategis', 'pelayanan', 'hotel', 'sangat', 'memuaskan']

3.3.2.8 Penghapusan *Stopwords*

Pada penelitian ini, tahapan penghapusan *stopwords* menggunakan *stopwords dictionary* dari pustaka NLTK. Perbandingan sebelum dan sesudah penghapusan *stopwords* dapat dilihat pada Tabel 3.6.

Tabel 3.6 Penghapusan *Stopwords*

Sebelum	Sesudah
['lokasi', 'pas', 'pelayanan', 'ramah']	['lokasi', 'pas', 'pelayanan', 'ramah']
['roomnya', 'sangat', 'nyaman', 'dan', 'bersih']	['roomnya', 'sangat', 'nyaman', 'bersih']
['lokasi', 'hotel', 'strategis', 'pelayanan', 'hotel', 'sangat', 'memuaskan']	['lokasi', 'strategis', 'pelayanan', 'sangat', 'memuaskan']

3.3.2.9 *Stemming*

Tahapan selanjutnya adalah melakukan *stemming* dengan menggunakan pustaka Sastrawi. Perbandingan sebelum dan sesudah *stemming* dapat dilihat pada Tabel 3.7.

Tabel 3.7 Stemming

Sebelum	Sesudah
['lokasi', 'pas', 'pelayanan', 'ramah']	'lokasi pas layan ramah'
['roomnya', 'sangat', 'nyaman', 'bersih']	'roomnya sangat nyaman bersih'
['lokasi', 'strategis', 'pelayanan', 'sangat', 'memuaskan']	'lokasi strategis layan sangat muas'

3.3.3 Anotasi Data

GPT (*Generative Pre-trained Transformer*) merupakan *large language models* yang memiliki keunggulan atau kemampuan dalam memahami dan berkomunikasi seperti manusia (Yenduri et al., 2023). Menurut Ding et al. (2022), GPT memiliki potensi untuk melakukan anotasi data dengan efisiensi lebih tinggi. Selain itu, menurut penelitian dari Nasution & Onan (2024) anotasi yang dihasilkan LLM memiliki kualitas yang kompetitif, terutama dalam analisis sentimen, anotasi manusia secara konsisten lebih unggul dalam tugas-tugas NLP yang lebih kompleks.

3.3.3.1 Zero-Shot Prompting

Dalam penelitian ini, proses anotasi memanfaatkan GPT sebagai anotator dengan menilai dan menganalisis ulasan ke dalam kelas positif atau negatif. Proses pelabelan data dilakukan menggunakan pendekatan *zero-shot classification* berbasis model GPT-4o-mini dari OpenAI. Menurut Y. Li (2023), *zero-shot* memiliki keunggulan, antara lain, muda diinterpretasikan, tidak memerlukan contoh, perancangan lebih sederhana, dan struktur *prompt* lebih fleksibel. Menurut Reynolds & McDonell (2021), *prompt zero-shot* yang dirancang dengan jelas dan tepat dapat mengungguli *few-shot* karena memberikan instruksi yang lebih

langsung dan tidak membingungkan model, sementara contoh dalam few-shot kadang disalahartikan sebagai bagian dari narasi.

3.3.4 Pemodelan Topik

Pemodelan topik dilakukan pada komentar dengan tujuan mengekstraksi aspek-aspek atau topik-topik utama yang terdapat di dalam korpus. Proses pemodelan topik dilakukan dengan memanfaatkan pustaka Python untuk pemodelan topik yang disebut Gensim. Proses pemodelan topik dimulai dengan membaca dataset yang selanjutnya akan dilakukan pra-pemrosesan data. Setelahnya akan dibuat daftar kata pada setiap komentar yang kemudian digunakan untuk membangun *dictionary*. Penelitian ini juga membangun model bigram dan trigram untuk membantu mengidentifikasi frasa dari kumpulan data. Tahap selanjutnya adalah membangun *dictionary*, kemudian dilanjutkan dengan membuat model LDA dengan TF-IDF.

3.3.5 Evaluasi Pemodelan Topik

Evaluasi model LDA dilakukan dengan menganalisis grafik *coherence score* dan *perplexity* menggunakan *elbow method* untuk menentukan model yang optimal. Model terpilih kemudian divisualisasikan dan dianalisis lebih lanjut. Analisis dilakukan dengan meneliti hubungan antar kata dalam setiap klaster yang dihasilkan oleh model LDA guna mengidentifikasi relevansi antar topik dalam berbagai klaster. Selain itu, penelitian ini memanfaatkan pustaka PyLDAvis sebagai alat bantu visualisasi untuk memahami topik yang dihasilkan. PyLDAvis menyajikan visualisasi dalam format .html, sehingga memudahkan eksplorasi keterkaitan antara topik dan istilah-istilah yang muncul di dalamnya. Setelah

evaluasi model dan pemilihan jumlah topik yang optimal, tahap selanjutnya adalah analisis topik terkait aspek pelayanan hotel.

3.3.6 K-Fold Cross Validation

Teknik validasi yang digunakan dalam penelitian ini adakah *K-Fold Cross Validation* dengan nilai $K = 5$. Menurut Berrar (2018), James et al. (2021), dan Raschka (2018), $k = 5$ pada *k-fold cross-validation* memberikan estimasi performa model yang stabil, akurat, dan efisien dalam mendeteksi *overfitting* atau *underfitting*.

3.3.7 Ekstraksi Fitur

Tahap ekstraksi fitur yang digunakan dalam penelitian ini adalah *words n-gram* dan TF-IDF, lebih tepatnya menggunakan kombinasi *words unigram*, *bigram*, *trigram*, *4-gram*, dan *5-gram*. Vektor dari fitur yang digunakan diekstraksi menggunakan TfidfVectorizer dari pustaka Scikit-Learn. Menurut Tripathy et al. (2016) dan Ibrohim & Budi (2019) kombinasi dari kelima fitur ini terbukti memberikan hasil kinerja model klasifikasi yang meningkat. Pada penelitian ini, tahap ekstraksi fitur hanya dijalankan pada analisis sentimen berbasis aspek.

Tabel 3.8 N-Grams

Sebelum	Sesudah
lokasi pas layan ramah	'lokasi', 'pas', 'layan', 'ramah', 'lokasi pas', 'pas layan', 'layan ramah', 'lokasi pas layan', 'pas layan ramah', 'lokasi pas layan ramah'
roomnya sangat nyaman bersih	'roomnya', 'sangat', 'nyaman', 'bersih', 'roomnya sangat', 'sangat nyaman', 'nyaman bersih', 'roomnya sangat nyaman', 'sangat nyaman bersih', 'roomnya sangat nyaman bersih'

Sebelum	Sesudah
lokasi strategis layanan sangat puas	'lokasi', 'layanan', 'sangat', 'strategis', 'puas', 'lokasi strategis', 'strategis layanan', 'layanan sangat', 'sangat puas', 'lokasi strategis layanan', 'strategis layanan sangat', 'layanan sangat puas', 'lokasi strategis layanan sangat', 'strategis layanan sangat puas', 'lokasi strategis layanan sangat puas'

Tabel 3.9 TF-IDF

Token	Term Count			Document Count	IDF	TF * IDF		
	Teks 1	Teks 2	Teks 3			Teks 1	Teks 2	Teks 3
bersih	0	0,25	0	1	0,477	0	0,119	0
layanan	0,25	0	0,2	2	0,176	0,044	0	0,035
layanan ramah	0,25	0	0	1	0,477	0,119	0	0
layanan sangat	0	0	0,2	1	0,477	0	0	0,095
layanan sangat puas	0	0	0,2	1	0,477	0	0	0,095
lokasi	0,25	0	0,2	2	0,176	0,044	0	0,035
lokasi pas	0,25	0	0	1	0,477	0,119	0	0
lokasi pas layanan	0,25	0	0	1	0,477	0,119	0	0
lokasi pas layanan ramah	0,25	0	0	1	0,477	0,119	0	0
lokasi strategis	0	0	0,2	1	0,477	0	0	0,095
lokasi strategis layanan	0	0	0,2	1	0,477	0	0	0,095
lokasi strategis layanan sangat	0	0	0,2	1	0,477	0	0	0,095

3.3.8 SMOTE

Dalam penelitian ini, jika terjadi ketidakseimbangan kelas dalam dataset maka akan dilakukan *Synthetic Minority Over-sampling Technique* atau SMOTE dari pustaka `imblearn.over_sampling`. Tahapan ini akan dilakukan setelah atau di dalam validasi silang. Menurut Muralidhar (2021), penggunaan SMOTE sesudah pembagian data menghasilkan skor validasi silang yang lebih akurat.

3.3.9 Analisis Sentimen Berbasis Aspek

Analisis sentimen berbasis aspek dilakukan pada komentar dengan mengklasifikasikan komentar berdasarkan aspek ke dalam kelas “positif” dan “negatif”. Penulis menggunakan dua model klasifikasi untuk mencari kinerja terbaik dalam klasifikasi sentimen dan aspek.

3.3.9.1 Naïve Bayes

Tahapan-tahapan model klasifikasi *Naïve Bayes* menggunakan data pada Tabel 3.8 sebagai data latih dan “*pelayanan ramah*” sebagai data uji, dijelaskan sebagai berikut.

1. Misalkan dokumen 1 pada data latih dikategorikan sebagai kelas A, dokumen 2 dikategorikan sebagai kelas B, dan dokumen 3 dikategorikan sebagai kelas A.
2. Data uji melalui tahap transformasi ekstraksi fitur berdasarkan Tabel 3.9.

Tabel 3.10 Transformasi Ekstraksi Fitur pada Data Uji

Sebelum	Sesudah
pelayanan ramah	‘layan’, ‘ramah’, layan ramah’

$$TF - IDF(layan) = \frac{1}{3} \times 0,176 = 0,058$$

$$TF - IDF(ramah) = \frac{1}{3} \times 0,477 = 0,159$$

$$TF - IDF(layan ramah) = \frac{1}{3} \times 0,477 = 0,159$$

3. Nilai probabilitas setiap kelas C pada semua data

$$P(A) = \frac{2}{3}$$

$$P(A) = \frac{1}{3}$$

4. Nilai rata-rata dan standar deviasi

Tabel 3.11 Rata-rata dan Standar Deviasi Kelas A

	layan	ramah	layan ramah
Teks 1	0,25	0,25	0,25
Teks 3	0,2	0	0
Mean	0,225	0,125	0,125
STD	0,035	0,177	0,177

Tabel 3.12 Rata-rata dan Standar Deviasi Kelas B

	layan	ramah	layan ramah
Teks 2	0	0	0

5. Fungsi densitas Gauss setiap atribut pada masing-masing kelas

$$f(\text{layan}|A) = \frac{1}{0,035\sqrt{2\pi}} e^{\frac{-(0,058-0,225)^2}{2(0,035)^2}} = 1,29763548 \times 10^{-4}$$

$$f(\text{layan, layan ramah}|A) = \frac{1}{0,177\sqrt{2\pi}} e^{\frac{-(0,159-0,125)^2}{2(0,177)^2}} = 2,213$$

$$f(\text{layan}|B) = \frac{1}{0\sqrt{2\pi}} e^{\frac{-(0,058-0)^2}{2(0)^2}} = \text{tidak terdefinisi}$$

$$f(\text{layan, layan ramah}|B) = \frac{1}{0\sqrt{2\pi}} e^{\frac{-(0,159-0)^2}{2(0)^2}} = \text{tidak terdefinisi}$$

6. Likelihood

$$P(X|A) = P(A) \times f(\text{layan}|A) \times f(\text{ramah}|A) \times f(\text{layan ramah}|A)$$

$$\begin{aligned} P(X|A) &= \frac{2}{3} \times 1,29763548 \times 10^{-4} \times 2,213 \times 2,213 \\ &= 4,23666652 \times 10^{-4} \end{aligned}$$

$$P(X|B) = P(B) \times f(\text{baik}|B) \times f(\text{ramah}|B) \times f(\text{baik ramah}|B)$$

$$P(X|B) = \text{tidak terdefinisi}$$

7. Keputusan prediksi adalah kelas A

3.3.9.2 Logistic Regression

Tahapan-tahapan model klasifikasi *Logistic Regression* menggunakan data pada Tabel 3.8 sebagai data latih dan “*pelayanan ramah*” sebagai data uji, dijelaskan sebagai berikut.

1. Misalkan dokumen 1 pada data latih dikategorikan sebagai kelas A, dokumen 2 dikategorikan sebagai kelas B, dan dokumen 3 dikategorikan sebagai kelas A.
2. Tahapan pelatihan model *logistic regression* melibatkan banyak perulangan guna menyesuaikan koefisien dan intersep agar meminimalkan fungsi loss. Pada kasus ini koefisien dan intersep diasumsikan sebagai berikut.

$$a = -0,71626337$$

Tabel 3.13 Asumsi Nilai Koefisien

Fitur	Bobot
bersih	0,21248827
layan	-0,21256607
layan ramah	-0,11387114
layan sangat	-0,09869494
layan sangat muas	-0,09869494
lokasi	-0,21256607
lokasi pas	-0,11387114
lokasi pas layan	-0,11387114
lokasi pas layan ramah	-0,11387114
lokasi strategis	-0,09869494
lokasi strategis layan	-0,09869494
lokasi strategis layan sangat	-0,09869494
lokasi strategis layan sangat muas	-0,09869494
muas	-0,09869494
nyaman	0,21248827
nyaman bersih	0,21248827
pas	-0,11387114
pas layan	-0,11387114
pas layan ramah	-0,11387114
ramah	-0,11387114
roomnya	0,21248827
roomnya sangat	0,21248827
roomnya sangat nyaman	0,21248827
roomnya sangat nyaman bersih	0,21248827
sangat	0,11379333
sangat muas	-0,09869494
sangat nyaman	0,21248827
sangat nyaman bersih	0,21248827
strategis	-0,09869494
strategis layan	-0,09869494
strategis layan sangat	-0,09869494
strategis layan sangat muas	-0,09869494

3. Menghitung nilai logit.

$$z = \alpha + \beta x$$

$$z = -0,71626337$$

$$+ ((0,21248827 \times 0) + (-0,21256607 \times 1))$$

$$+ (-0,11387114 \times 1) + \dots + (-0,09869494 \times 0))$$

$$z = -0,71626337 + (-0,44030835) = -1,1565172$$

4. Menghitung probabilitas menggunakan persamaan regresi logistik atau fungsi sigmoid.

$$\pi(x) = \frac{1}{1 + e^{-(-1,1565172)}} = 0,23929078$$

5. Menentukan kelas data uji, dengan ketentuan sebagai berikut.

$$Prediksi = \begin{cases} 0, & \text{jika } \pi(x) \leq 0,5 \\ 1, & \text{jika } \pi(x) > 0,5 \end{cases}$$

6. Berdasarkan poin 5, data uji diklasifikasikan sebagai kelas 0 atau kelas A.

3.3.10 Evaluasi

Proses evaluasi model pembelajaran mesin menggunakan akurasi yang perhitungannya terdapat pada persamaan. Model terbaik dipilih berdasarkan nilai rata-rata akurasi dari keseluruhan *fold*.

BAB IV

HASIL DAN PEMBAHASAN

4.1 Pengumpulan Data

Dari proses pengumpulan data, terkumpul sebanyak 6.430 ulasan yang kemudian akan dilanjutkan ke tahap pra-pemrosesan. Sebagian besar ulasan tersebut ditulis dalam bahasa Indonesia. Adapun tampilan dari data yang diperoleh dapat dilihat pada Gambar 4.1.

	name	rating	duration	review
0	Brigita Sinaga	1.0	a month ago	pulang ke hotel jam setengah 12 malam, mendapa...
1	Yusika	2.0	a month ago	Kamar hotel bagus tp kurang bersih. Kolam rena...
2	shilfia yuniar	4.0	a month ago	Hotel strategis ada di pusat kota. Sayangnya w...
3	Rizki Gutama	5.0	a month ago	Hotel yang lokasinya cukup strategis di tengah...
4	ukhtiy affiah	1.0	a month ago	Buruk.Baru nginep pertama kali.Sudah telpon ak...
...	...	...	...	...
6425	Sahida Nafisah	4.0	a year ago	Nearby activities: ribut
6426	WINDY RAHAYU HOETARMAN	5.0	2 years ago	Grand Zuri is solid
6427	Said Reza Ilham Pratama	5.0	3 years ago	Mantaaap
6428	Andi Baso Asrul	5.0	6 years ago	Hotel Grand Zuri is great
6429	Dimdims Andhika	5.0	5 years ago	Good hotel in Pku

6430 rows × 4 columns

Gambar 4.1 Hasil Pengumpulan Data

4.2 Pra-pemrosesan Data

Tahapan pra-pemrosesan data meliputi beberapa tahapan, antara lain *case folding*, penghapusan emoji, angka, tanda baca, *whitespace*, dan karakter tunggal, penghapusan kata tunggal, penghapusan kalimat berbahasa Inggris, tokenisasi, penghapusan *stopwords*, dan *stemming*. Tahapan ini memperoleh 5178 ulasan

bersih yang siap digunakan pada tahapan berikutnya. Adapun beberapa data yang telah melalui tahapan pra-pemrosesan ditampilkan pada Tabel 4.1.

Tabel 4.1 Sebelum dan Sesudah Pra-pemrosesan

No	Sebelum	Sesudah
1.	pulang ke hotel jam setengah 12 malam, mendapati kamar bocor dan udah mulai sangat becek, tapi penanganannya sangat lambat. untuk sampai engineer-nya datang ke kamar buat ngecek aja harus sampai ditelfon 2 kali. pas udah dicek engineer-nya masih berusaha buat benerin kebocoran dan tidak memberikan solusi untuk ganti kamar, padahal kondisi kamar udah becek...	pulang jam malam dapat kamar bocor sangat becek tangan sangat lambat engineer-nya kamar ngecek ditelfon kali pas cek engineer-nya usaha benerin bocor tidak solusi ganti kamar kondisi kamar becek...
2.	Kamar hotel bagus tp kurang bersih. Kolam renang-nya jorok banget, berlumut sepertinya tidak pernah dibersihkan. Seumur2 baru ini hotel yg hanya menyediakan 1 handuk untuk yg berenang. (kebanyakan 2/kamar)...	kamar bagus kurang bersih kolam renang jorok lumut tidak bersih umur sedia handuk renang banyak kamar...
3.	Hotel strategis ada di pusat kota. Sayangnya waktu cekin agak lama ngantri waktu itu. Sarapan bervariasi & enak. Kamar non smoking tapi dapat sarung bantal dgn bekas sundutan rokok. Utk kolam renang-nya kotor, keramik licin berlumut	strategis pusat kota sayang cekin agak ngantri sarap variasi enak kamar non smoking sarung bantal bekas sundut rokok kolam renang kotor keramik licin lumut
4.	Hotel yang lokasinya cukup strategis di tengah kota, ada jogging track nya, pelayanannya dr receptionist sama petugas di restonya bagus semua, tipikal kamar hotel lama yang lantainya ditutup karpet tapi bersih semua. Layanan laundrynya juga bersih, cepet. Minusnya cuma klosetnya saat ngeflush harus ditahan	lokasi cukup strategis tengah kota jogging track layan receptionist tugas restonya bagus tipikal kamar lantai tutup karpet bersih layan laundrynya bersih cepat minus kloset ngeflush tahan
5.	Buruk. Baru ngeinap pertama kali. Sudah telpon akan late check in	buruk inap kali telepon late check kamar non smoking confirm kamar

No	Sebelum	Sesudah
	minta kamar non smoking, confirm ada 2 kamar.Pas datang gak ada sama sekali. TRUS APA GUNANYA DITELPON CONFIRMED BISA???Staff kurang ramah, suaranya diperbesar, diperjelas kalau ngomong...	pas tidak guna telpon confirmed staff kurang ramah suara besar jelas bicara...

4.3 Anotasi Data

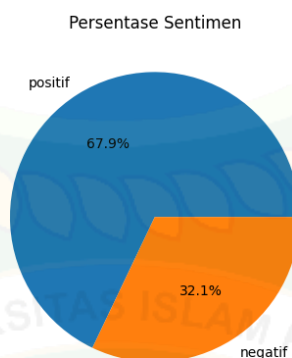
Setelah tahapan pra-pemrosesan data, selanjutnya data akan melalui tahapan anotasi sentimen. Tahapan ini dilakukan secara otomatis menggunakan model *gpt-4o-mini* dengan teknik *zero-shot prompting* dengan pesan *prompt* “Anda adalah model klasifikasi teks yang ahli dalam menganalisis ulasan hotel dan menentukan apakah ulasan tersebut bersifat positif atau negatif. Jawablah hanya dengan salah satu dari dua label itu, tanpa penjelasan.”. Hasil dari tahapan ini ditampilkan dalam bentuk tabel pada Tabel 4.2. Dari Tabel 4.2, terdapat 8 ulasan yang tidak teranotasi dengan baik oleh sistem. Adapun perbandingan sentimen dalam bentuk persentase dapat dilihat pada **Error! Reference source not found..**

Tabel 4.2 Hasil Anotasi Sentimen

Sentimen	Jumlah
Positif	3511
Negatif	1659
Total	5170

Berdasarkan Tabel 4.2, hasil anotasi sentimen menunjukkan bahwa dari total 5.170 data ulasan yang digunakan dalam penelitian, sebanyak 3.511 ulasan

diklasifikasikan sebagai sentimen positif dan 1.659 ulasan diklasifikasikan sebagai sentimen negatif.



Gambar 4.2 Persentase Sentimen

4.4 Pemodelan Topik

Langkah awal dalam proses pemodelan topik dimulai dengan membangun daftar kata, kemudian dilanjutkan dengan pembuatan *dictionary* serta pembangunan model TF-IDF. Setelah itu, model LDA dibuat dan dievaluasi berdasarkan 2 metrik, yaitu metrik *coherence* dan *perplexity*. Kedua metrik ini merupakan landasan dalam analisis pemilihan aspek yang akan digunakan pada tahapan analisis sentimen berbasis aspek. Tahapan ini diakhiri dengan pemilihan model terbaik dengan batasan jumlah topik sebanyak 10.

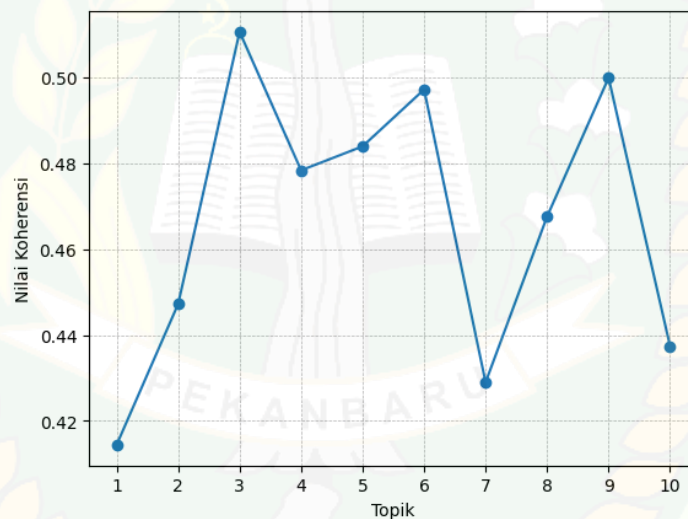
4.5 Elbow Method

Pemilihan model topik terbaik dilakukan dengan menganalisis grafik *coherence* dan *perplexity*, di mana model paling optimal ditentukan melalui metode analisis *elbow*. Adapun nilai dari metrik *coherence* dan *perplexity* dapat dilihat pada Tabel 4.3.

Tabel 4.3 Nilai *Coherence* dan *Perplexity*

Jumlah Topik	<i>Coherence Value</i>	<i>Perplexity Value</i>
1	0.414417	-6.380.032
2	0.447212	-6.593.265
3	0.510619	-6.733.134
4	0.478368	-6.918.059
5	0.483953	-7.038.549
6	0.497192	-7.184.044
7	0.428990	-7.271.644
8	0.467603	-7.364.776
9	0.499914	-7.449.626
10	0.437405	-7.575.381

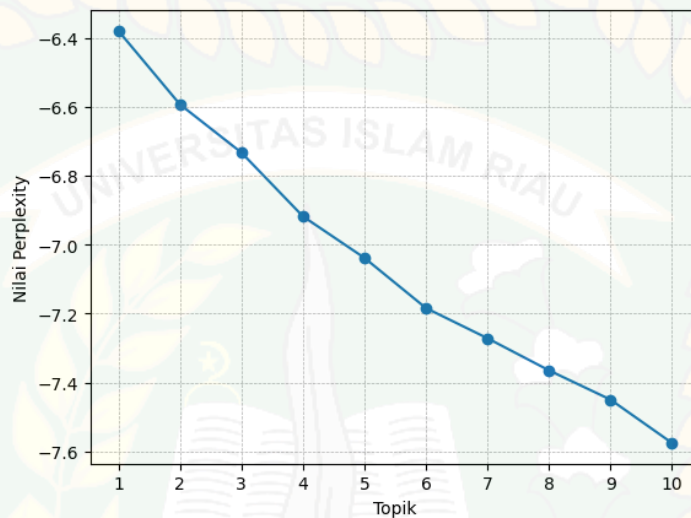
4.5.1 *Coherence Score*

**Gambar 4.3** Perbandingan *Coherence Value*

Dari Gambar 4.3, terlihat bahwa nilai *coherence* mengalami fluktuasi seiring bertambahnya jumlah topik. Nilai *coherence* tertinggi dicapai pada topik ke-3, yang menunjukkan bahwa model dengan tiga topik menghasilkan pembagian topik yang paling koheren dan bermakna dibandingkan jumlah topik lainnya. Setelah topik ke-3, nilai *coherence* cenderung naik-turun, dengan sedikit peningkatan pada topik ke-

6 dan ke-9, namun tidak melebihi nilai maksimum pada topik ke-3. Dengan demikian, model dengan tiga topik dapat dianggap sebagai pilihan terbaik untuk menggambarkan struktur topik dalam data ini.

4.5.2 *Perplexity*



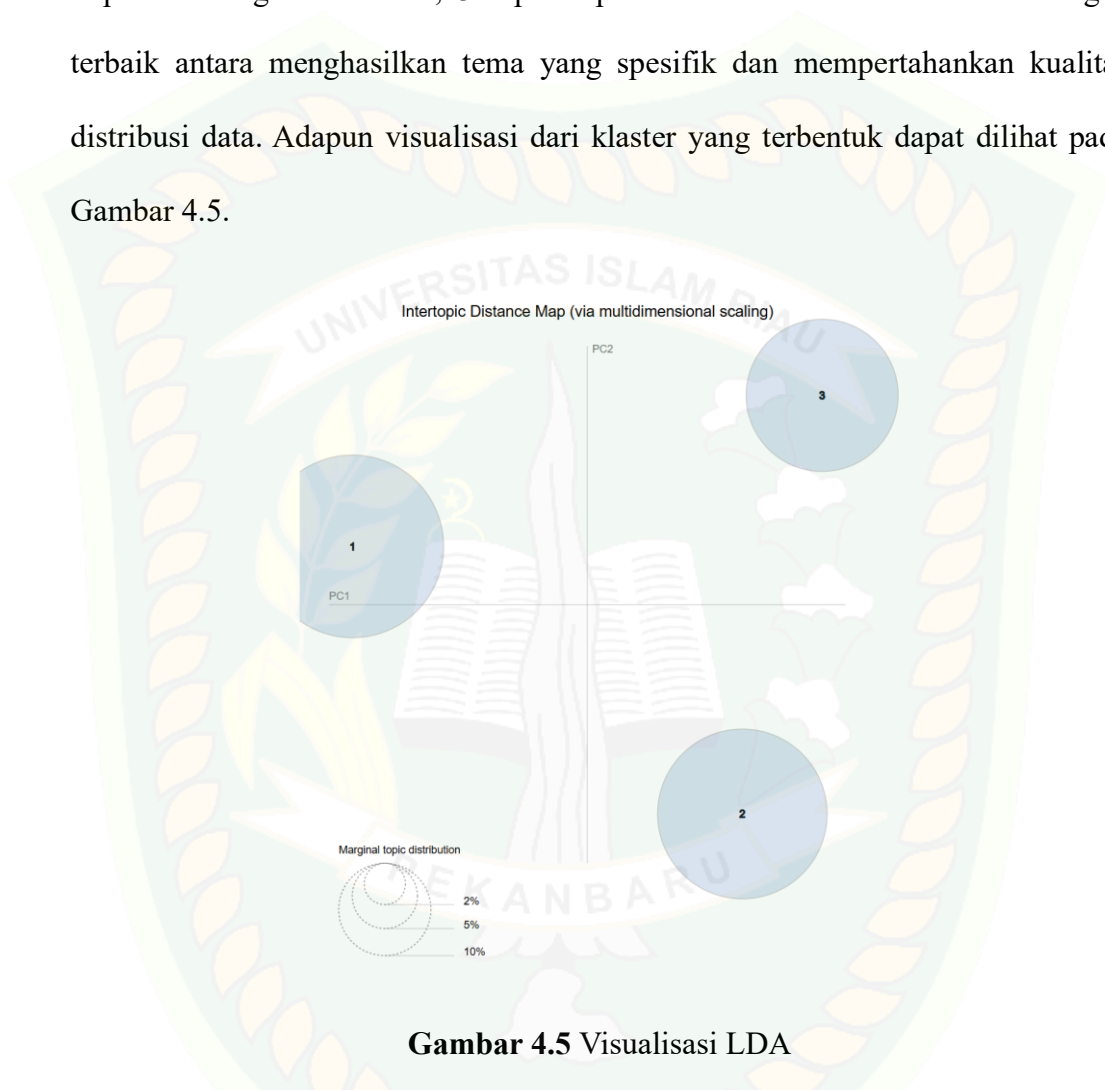
Gambar 4.4 Perbandingan *Perplexity Value*

Berdasarkan Gambar 4.4, terlihat bahwa seiring dengan peningkatan jumlah topik dari 1 hingga 10, nilai *perplexity* menunjukkan kecenderungan menurun. Penurunan ini mengindikasikan bahwa model secara keseluruhan menjadi lebih baik atau lebih mampu memprediksi sampel. Namun, penurunan yang paling signifikan terjadi pada topik-topik awal dan laju penurunannya mulai melandai setelah mencapai sekitar 4 atau 5 topik.

4.6 Penentuan Jumlah Topik

Berdasarkan uraian sebelumnya, jumlah topik optimal berada pada rentang 3 hingga 5 topik. Model 3 topik dipilih untuk langkah selanjutnya berdasarkan keseimbangan *coherence* dan *perplexity*. Pemilihan ini dikarenakan 3 topik

menghasilkan nilai *coherence* tertinggi di seluruh pengujian, dan pada titik ini, nilai *perplexity* sudah menunjukkan penurunan tajam sebelum mulai mendatar setelah Topik 4. Dengan demikian, 3 topik dipilih karena memberikan keseimbangan terbaik antara menghasilkan tema yang spesifik dan mempertahankan kualitas distribusi data. Adapun visualisasi dari kluster yang terbentuk dapat dilihat pada Gambar 4.5.



Dari Gambar 4.5, dapat dilihat bahwa posisi dari ketiga kluster terpisah satu sama lain. Hal ini menunjukkan bahwa setiap topik memiliki perbedaan tema yang cukup jelas dan tidak saling tumpang tindih secara signifikan. Adapun kata kunci teratas dari model yang terbentuk adalah sebagai berikut.

Tabel 4.4 Kata Kunci Model Pemodelan Topik

Topik	Probabilitas Kata
1	0.024*"tidak" + 0.018*"kurang" + 0.015*"air" + 0.011*"mandi" + 0.010*"parkir" + 0.009*"agak" + 0.009*"conditioner" + 0.009*"lumayan" + 0.008*"lantai" + 0.007*"dingin"
2	0.031*"layan" + 0.030*"ramah" + 0.028*"bagus" + 0.025*"enak" + 0.024*"sangat" + 0.020*"bersih" + 0.018*"makan" + 0.017*"nyaman" + 0.014*"rooms" + 0.013*"oke"
3	0.022*"nyaman" + 0.020*"kota" + 0.020*"strategis" + 0.020*"sangat" + 0.019*"dekat" + 0.019*"lokasi" + 0.016*"bersih" + 0.015*"pusat" + 0.014*"bagus" + 0.013*"fasilitas"

Berdasarkan Tabel 4.4, terbentuk tiga topik utama dengan karakteristik kata kunci yang berbeda. Adapun uraian dalam penentuan aspek berdasarkan Tabel 4.4 dijelaskan sebagai berikut:

1. Topik 1: Topik ini membahas mengenai fasilitas yang tersedia di hotel, seperti *air*, *mandi*, *parkir*, *conditioner*, dan *lantai*. Kata-kata tersebut menunjukkan bahwa aspek yang dibahas berkaitan dengan ketersediaan serta kondisi sarana dan prasarana hotel yang digunakan oleh tamu selama menginap.
2. Topik 2: Topik ini berkaitan dengan aspek pelayanan dan kenyamanan tamu selama berada di hotel. Kata-kata seperti *layan*, *ramah*, *bagus*, *enak*, *bersih*, dan *nyaman* menggambarkan hal-hal yang berhubungan dengan interaksi staf hotel, kebersihan, dan kualitas layanan.
3. Topik 3: Topik ini membahas mengenai lokasi. Kata-kata seperti *kota*, *strategis*, *dekat*, *lokasi*, dan *pusat* menunjukkan bahwa pembahasan dalam topik ini berfokus pada letak geografis hotel serta kedekatannya dengan pusat kota atau fasilitas umum lainnya.

Dari uraian di atas, dapat disimpulkan bahwa topik 1 merupakan topik yang membahas mengenai aspek fasilitas, topik 2 membahas mengenai aspek pelayanan, dan topik 3 membahas mengenai aspek lokasi. Sehingga aspek yang digunakan pada tahap analisis sentimen berbasis aspek adalah aspek fasilitas, pelayanan, dan lokasi. Adapun pengelompokan kata-kata ke dalam topik-topik tertentu berdasarkan pola kemunculan dan hubungan antar kata dalam teks. Dapat dilihat pada Tabel 4.5, Sentimen positif aspek lokasi memperoleh jumlah tertinggi dengan 1.600 ulasan, aspek pelayanan sebanyak 1.263 ulasan, dan aspek fasilitas sebanyak 648 ulasan. Sementara pada sentimen negatif, aspek fasilitas memperoleh jumlah paling tinggi dengan 1.348 ulasan yang mengindikasikan bahwa fasilitas menjadi aspek yang paling banyak mendapat keluhan.

Tabel 4.5 Distribusi Aspek

Sentimen/ Aspek	Fasilitas	Pelayanan	Lokasi	Jumlah
Positif	648	1263	1600	3511
Negatif	1348	137	174	1659
Jumlah	1996	1400	1774	5170

4.7 Analisis Sentimen Berbasis Aspek

Analisis sentimen berbasis aspek dilakukan dengan membandingkan kinerja dua algoritma pembelajaran mesin, yaitu *Naïve Bayes* dan *Logistic Regression*. Analisis sentimen berbasis aspek dilakukan menjadi dua tahapan, tahapan pertama berfokus untuk mengklasifikasikan aspek dan tahapan kedua berfokus untuk mengklasifikasikan sentimen berdasarkan aspek. Sebelum model dilatih, data akan melalui tahapan ekstraksi fitur menggunakan metode *n-grams* dan TF-IDF agar informasi teks dapat diolah secara optimal oleh model.

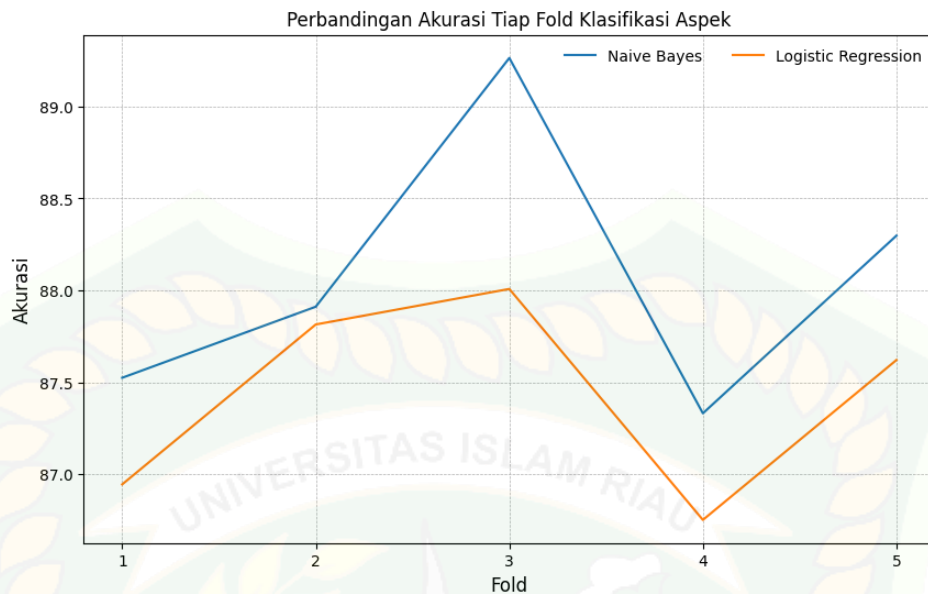
Pada proses pelatihan model, ditemukan bahwa distribusi data pada setiap aspek tidak seimbang, sehingga diterapkan metode SMOTE untuk menambah jumlah data pada kelas dengan sampel lebih sedikit. Proses pelatihan ini dijalankan menggunakan *pipeline* dari pustaka scikit-learn guna mempermudah alur pemodelan. Selanjutnya model dilatih menggunakan validasi silang dengan skema *k-fold* sebanyak 5 *folds* untuk memastikan performa model terukur secara konsisten.

Validasi silang menghasilkan nilai akurasi untuk setiap *fold* yang digunakan dalam pengujian model. Perbandingan akurasi dari masing-masing *fold* ditunjukkan pada Tabel 4.6.

Tabel 4.6 Rata-rata Akurasi Model Klasifikasi Aspek

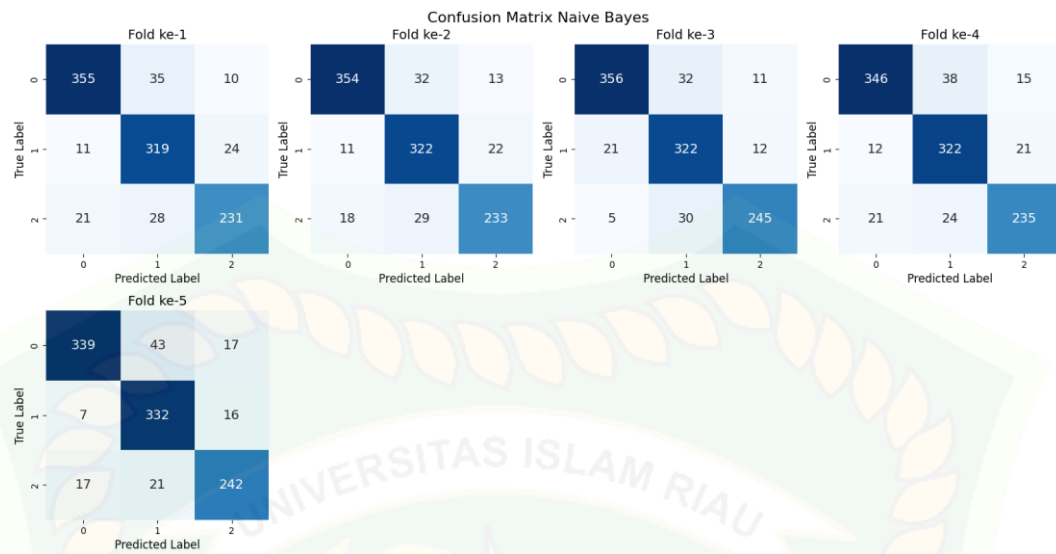
Model/Fold	1	2	3	4	5	Mean
Naive Bayes	87.5%	87.9%	89.3%	87.3%	88.3%	88.1%
Logistic Regression	86.9%	87.8%	88.0%	86.8%	87.6%	87.4%

Berdasarkan tabel di atas, model *Naïve Bayes* menunjukkan nilai akurasi tertinggi dibandingkan model lainnya. Adapun visualisasi dari Tabel 4.6 adalah sebagai berikut.

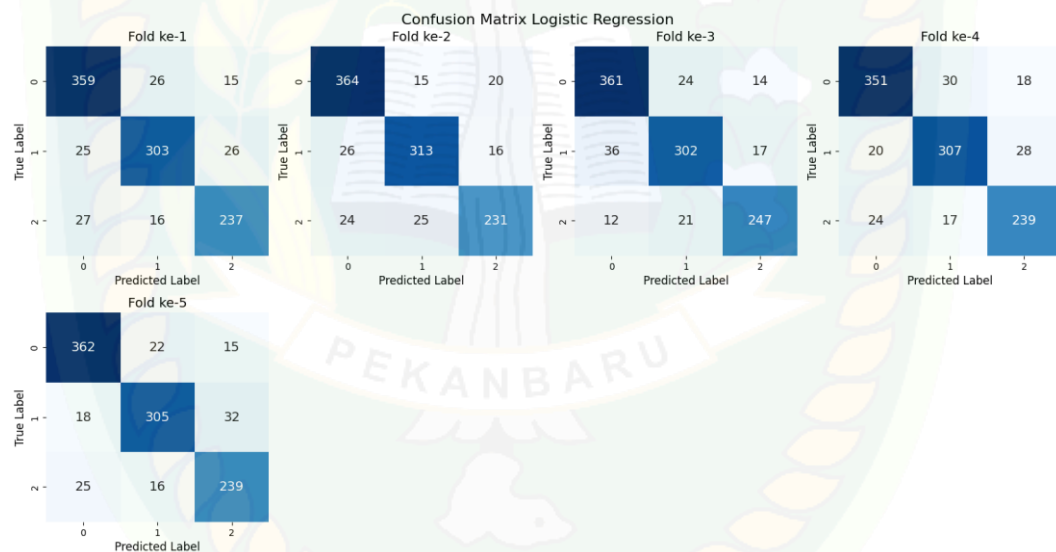


Gambar 4.6 Perbandingan Akurasi Tiap *Fold* Klasifikasi Aspek

Pada Gambar 4.6, grafik memperlihatkan perbandingan akurasi model *Naive Bayes* dan *Logistic Regression* pada lima fold pengujian. *Naive Bayes* menunjukkan performa yang lebih baik dibandingkan *Logistic Regression* hampir semua fold. *Naive Bayes* mencapai akurasi tertinggi pada fold ketiga yaitu 89,3%, sedangkan *Logistic Regression* hanya mencapai akurasi tertingginya pada fold yang sama dengan nilai 88%. Meskipun kedua model mengalami penurunan akurasi pada fold keempat, *Naive Bayes* tetap mempertahankan nilai akurasi yang lebih tinggi. Hal ini menunjukkan *Naive Bayes* lebih unggul dan lebih stabil dalam melakukan klasifikasi aspek. Selanjutnya *confussion matrix* dari model *Naive Bayes* dan *Logistic Regression* pada klasifikasi aspek yang ditampilkan pada Gambar 4.7 dan Gambar 4.8.



Gambar 4.7 *Confussion Matrix* Model *Naïve Bayes* Pada Klasifikasi Aspek



Gambar 4.8 *Confussion Matrix* Model *Logistic Regression* Pada Klasifikasi Aspek

Berdasarkan hasil *Confussion Matrix* pada lima fold, kedua model menunjukkan kemampuan klasifikasi yang baik, ditunjukkan oleh dominasi nilai diagonal pada setiap fold. Model *Naïve Bayes* terlihat lebih stabil terutama pada kelas 1 dan kelas 2 yang secara konsisten memiliki jumlah prediksi benar yang tinggi. Meskipun

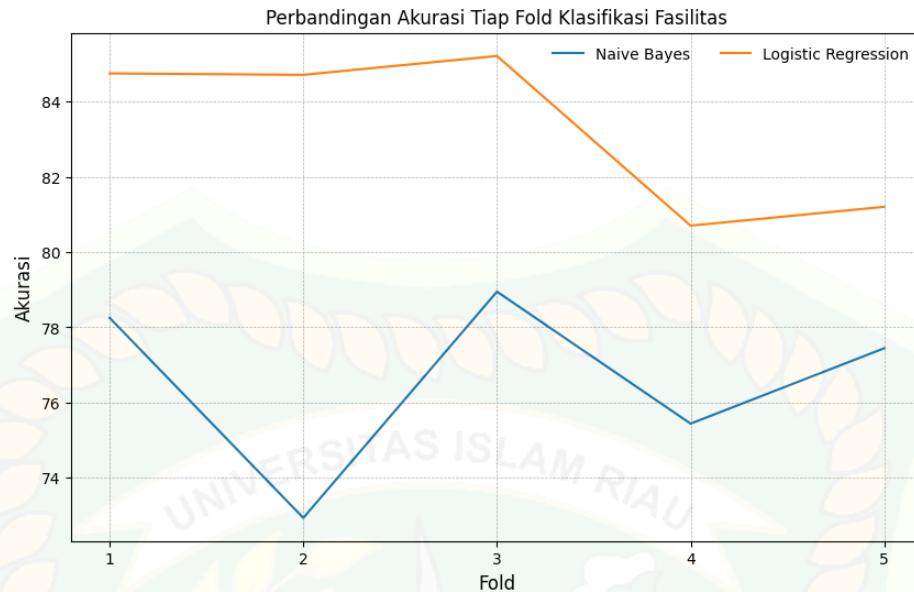
kesalahan klasifikasi antar kelas tetap terjadi jumlahnya relatif kecil dan tidak memengaruhi konsistensi performa model secara keseluruhan. Sedangkan *Logistic Regression* juga menunjukkan pola prediksi yang konsisten tetapi jumlah prediksi benarnya pada setiap fold sedikit lebih rendah dibandingkan *Naive Bayes* terutama pada kelas 1 dan kelas 2. Hasil ini mengindikasikan bahwa *Naive Bayes* memiliki performa yang lebih baik dan lebih stabil dibandingkan *Logistic Regression* dalam proses klasifikasi pada pengujian *k-fold* yang dilakukan.

Langkah berikutnya yaitu melakukan pelatihan model klasifikasi sentimen berdasarkan aspek. Pada tahap ini, model dikembangkan dan dilatih menggunakan data yang telah dipisahkan sesuai dengan masing-masing aspek, sehingga model dapat mempelajari pola sentimen secara spesifik pada tiap aspek yang dianalisis. Hasil dari validasi silang untuk kedua aspek ditampilkan pada Tabel 4.7.

Tabel 4.7 Rata-rata Akurasi Model Klasifikasi Sentimen

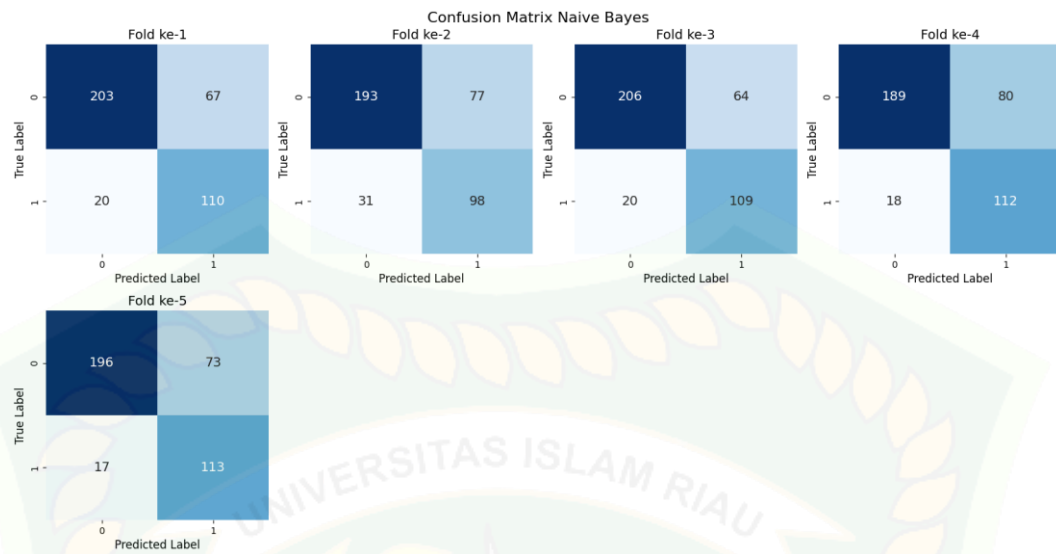
Aspek	Model/Fold	1	2	3	4	5	Mean
Fasilitas	Naive Bayes	78.2%	72.9%	78.9%	75.4%	77.4%	76.6%
	Logistic Regression	84.8%	84.7%	85.2%	80.7%	81.2%	83.3%
Pelayanan	Naive Bayes	93.6%	90.4%	92.5%	90.7%	92.1%	91.9%
	Logistic Regression	95.0%	91.4%	94.6%	94.3%	92.9%	93.6%
Lokasi	Naive Bayes	87.0%	85.9%	88.2%	87.6%	88.7%	87.5%
	Logistic Regression	90.1%	89.9%	91.5%	91.0%	93.5%	91.2%

Berdasarkan Tabel 4.7, dapat disimpulkan bahwa *Logistic Regression* lebih efektif dalam mengklasifikasikan sentimen berdasarkan aspek. Adapun akurasi dari tiap *fold* atau lipatan hasil validasi silang pada aspek fasilitas dapat dilihat sebagai berikut.

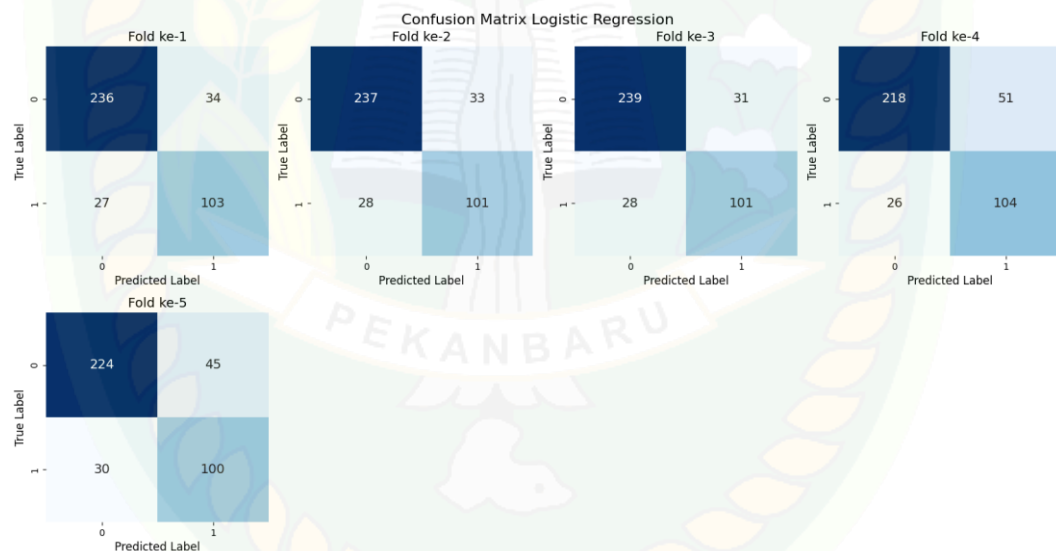


Gambar 4.9 Perbandingan Akurasi Tiap *Fold* Klasifikasi Sentimen Aspek Fasilitas

Gambar 4.9 menunjukkan perbandingan akurasi model *Naive Bayes* dan *Logistic Regression* pada lima fold untuk klasifikasi sentimen aspek fasilitas. Model *Logistic Regression* memperoleh akurasi yang lebih tinggi dan lebih stabil dibandingkan *Naive Bayes* pada seluruh fold. *Logistic Regression* mempertahankan akurasi di kisaran 81-85% dengan nilai tertinggi pada fold ketiga. *Naive Bayes* menunjukkan ketidakstabilan yang cukup signifikan dengan nilai terendah pada fold kedua dan meningkat kembali pada fold ketiga dan kelima. Hasil ini menunjukkan bahwa *Logistic Regression* lebih unggul dan konsisten dalam melakukan klasifikasi sentimen pada aspek fasilitas dibandingkan *Naive Bayes*. Adapun *confussion matrix* klasifikasi sentimen pada aspek fasilitas dari model *Naive Bayes* dan *Logistic Regression* ditampilkan pada Gambar 4.10 dan Gambar 4.11.



Gambar 4.10 *Confusion Matrix Model Naïve Bayes* pada Klasifikasi Sentimen Aspek Fasilitas

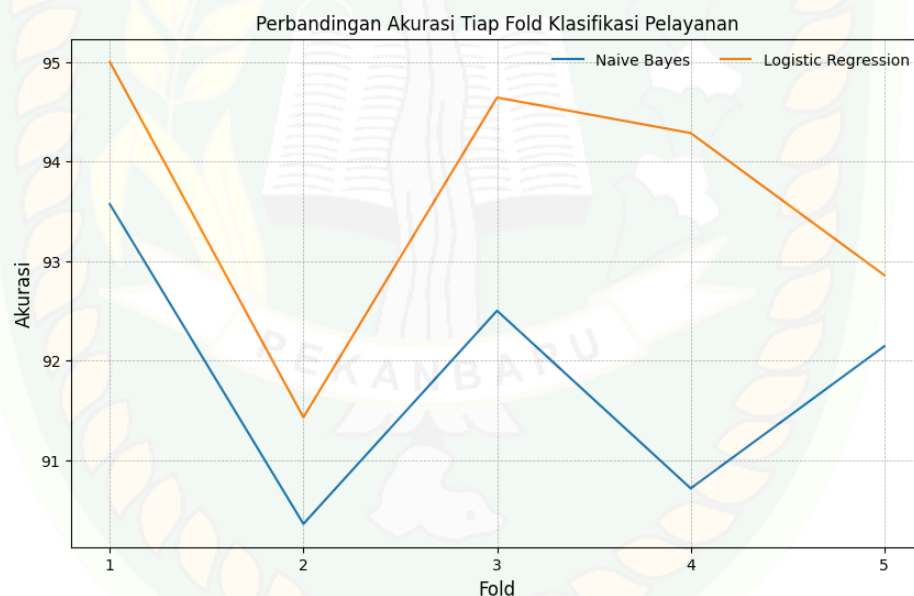


Gambar 4.11 *Confusion Matrix Model Logistic Regression* pada Klasifikasi Sentimen Aspek Fasilitas

Hasil evaluasi melalui 5-fold *cross validation* menunjukkan bahwa *Naive Bayes* mampu mengklasifikasikan kelas 0 dengan cukup baik ditandai oleh nilai *true negative* yang konsisten pada setiap fold. Namun model ini masih

menghasilkan jumlah *false positive* dan *false negative* yang relatif tinggi sehingga ketepatan dalam membedakan kelas belum optimal. Sedangkan *Logistic Regression* menunjukknn perfoma yang lebih baik dan stabil dibandingkan *Naive Bayes*. Hal ini terlihat dari nilai *true positive* dan *true negative* yang lebih tinggi serta jumlah *false positive* dan *false negative* yang lebih rendah pada seluruh fold. Konsistensi hasil ini mengindikasikan bahwa *Logistic Regression* mampu memodelkan hubungan antar fitur dengan lebih akurat sehingga menghasilkan prediksi yang lebih tepat.

Perbandingan akurasi dari tiap *fold* atau lipatan hasil validasi silang pada aspek pelayanan dapat dilihat sebagai berikut.

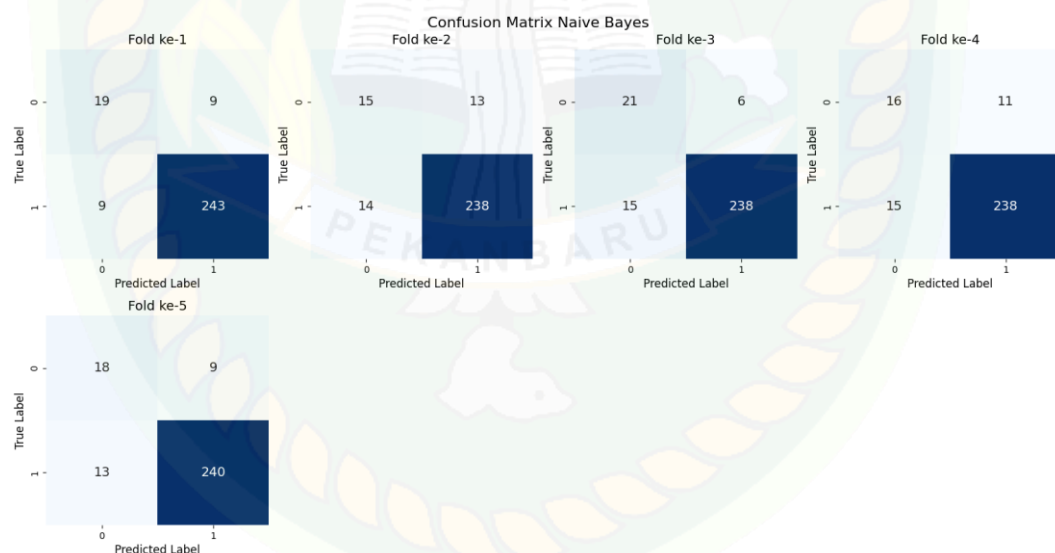


Gambar 4.12 Perbandingan Akurasi Tiap *Fold* Klasifikasi Sentimen Aspek Pelayanan

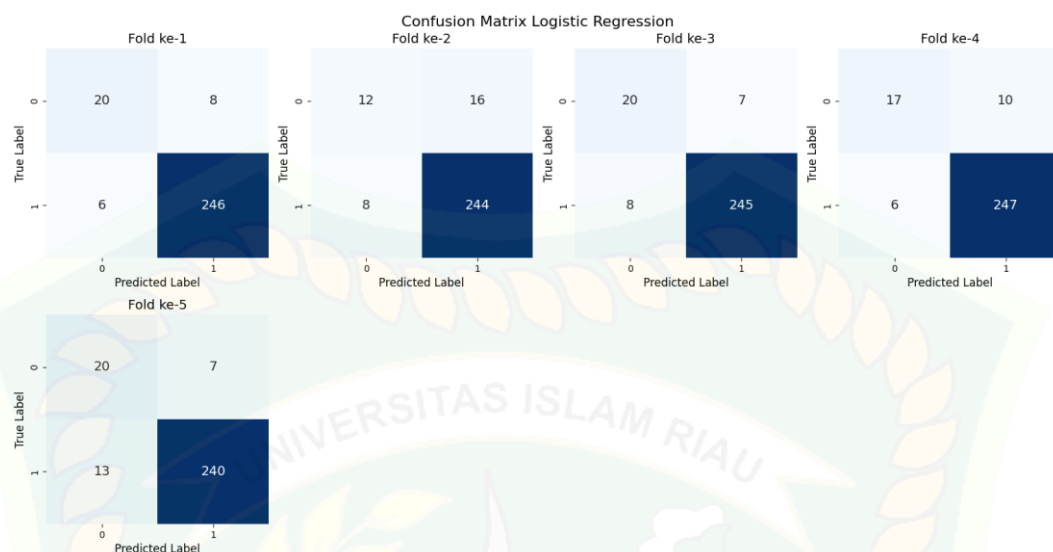
Berdasarkan Gambar 4.12 menunjukkan perbandingan akurasi model *Naive Bayes* dan *Logistic Regression* pada setiap fold proses *cross validation* untuk klasifikasi sentimen aspek pelayanan. Dari grafik terlihat bahwa *Logistic*

Regression secara konsisten menghasilkan akurasi yang lebih tinggi dibandingkan *Naive Bayes* pada seluruh fold. Model *Logistic regression* mencapai akurasi tertinggi pada fold pertama dan ketiga dengan nilai mendekati 95% sedangkan penurunan akurasi terjadi pada fold kelima namun tetap berada diatas 93%. Sedangkan model *Naive Bayes* menunjukkan ketidakstabilan akurasi yang lebih tajam. Akurasi tertinggi dicapai pada fold pertama, kemudian mengalami penurunan signifikan pada fold kedua sebelum kembali meningkat pada fold ketiga. Meskipun terjadi perbaikan akurasi *Naive Bayes* tetap berada di bawah *Logistic Regression*.

Confussion matrix klasifikasi sentimen pada aspek fasilitas dari model *Naïve Bayes* dan *Logistic Regression* ditampilkan pada Gambar 4.13 dan Gambar 4.14.



Gambar 4.13 *Confusion Matrix* Model *Naïve Bayes* pada Klasifikasi Sentimen Aspek Pelayanan

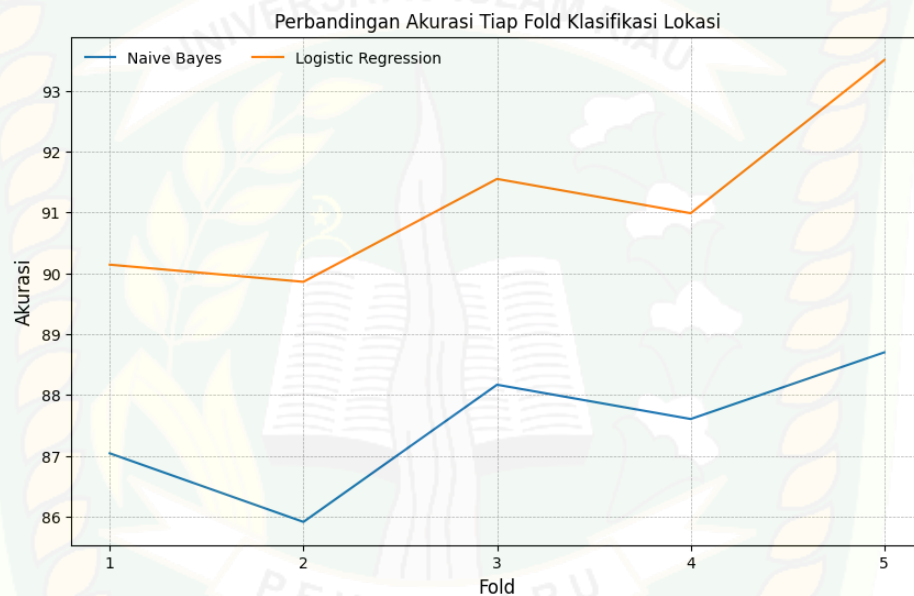


Gambar 4.14 *Confusion Matrix Model Logistic Regression* pada Klasifikasi Sentimen Aspek Pelayanan

Gambar 4.13 menunjukkan bahwa model memiliki kemampuan klasifikasi yang baik pada kelas 1, yang ditunjukkan oleh nilai *true positive* yang konsisten tinggi pada setiap fold, yaitu berada pada kisaran 238 hingga 243. Namun demikian, performa pada kelas 0 masih kurang optimal. Hal ini terlihat dari nilai *true negative* yang relatif rendah serta jumlah *false positive* dan *false negative* yang masih muncul pada seluruh fold. Kondisi tersebut mengindikasikan bahwa *Naive Bayes* lebih mampu mengenali kelas mayoritas, tetapi belum efektif dalam membedakan kelas minoritas. Sedangkan model *Logistic Regression* menunjukkan performa yang lebih akurat dan stabil dibandingkan *Naive Bayes*. Nilai *true positive* yang dihasilkan berada pada rentang 240 hingga 247, dengan nilai *true negative* yang lebih tinggi dan konsisten di seluruh fold. Jumlah *false positive* dan *false negative* juga relatif lebih rendah, yang menunjukkan bahwa *Logistic Regression* memiliki

kemampuan pemisahan kelas yang lebih baik. Secara keseluruhan, hasil ini mengonfirmasi bahwa Logistic Regression lebih unggul dalam mengklasifikasikan sentimen aspek pelayanan, dengan tingkat kesalahan prediksi yang lebih rendah pada setiap fold.

Perbandingan akurasi dari tiap *fold* atau lipatan hasil validasi silang pada aspek pelayanan dapat dilihat sebagai berikut.

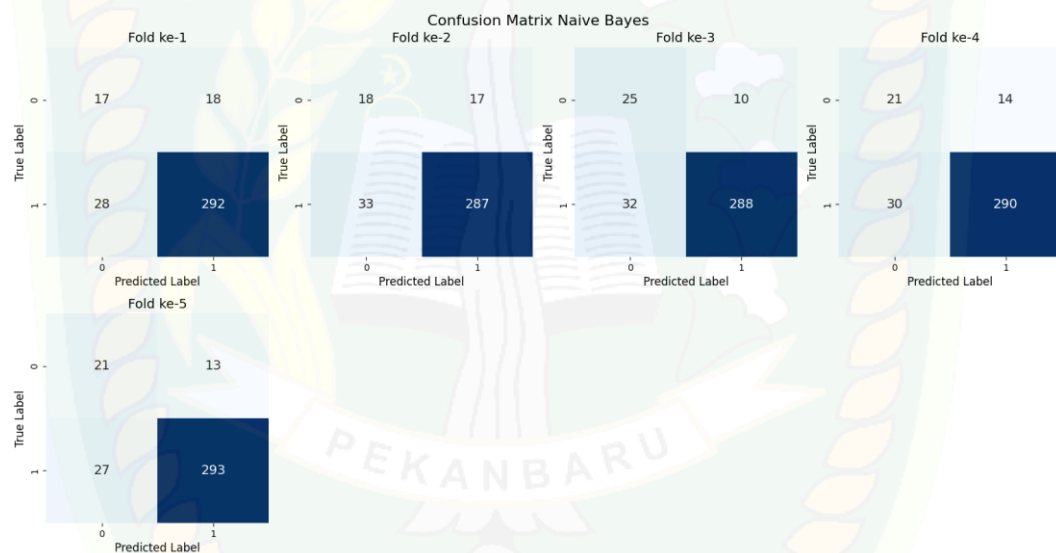


Gambar 4.15 Perbandingan Akurasi Tiap *Fold* Klasifikasi Sentimen Aspek Lokasi

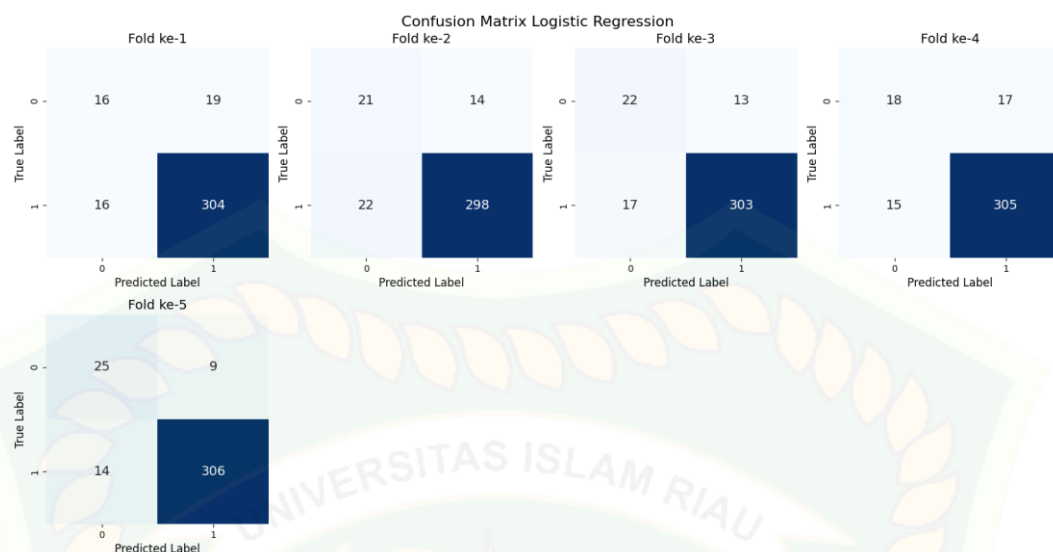
Gambar 4.15 menunjukkan perbandingan akurasi model Naive Bayes dan Logistic Regression pada setiap fold proses *cross validation* untuk klasifikasi sentimen aspek lokasi. Berdasarkan grafik tersebut, terlihat bahwa model *Logistic Regression* secara konsisten menghasilkan akurasi yang lebih tinggi dibandingkan *Naive Bayes* pada seluruh fold. *Logistic Regression* mencapai akurasi tertinggi pada fold kelima dengan nilai sekitar 94%, sedangkan akurasi terendah terjadi pada fold kedua, namun tetap berada di atas 89%. Hal ini menunjukkan bahwa performa

model relatif stabil dan tidak mengalami penurunan signifikan pada setiap pengujian. Sedangkan model *Naive Bayes* menunjukkan akurasi yang lebih rendah dan mengalami ketidakstabilan yang lebih besar pada tiap fold. Akurasi tertinggi dicapai pada fold kelima, yaitu sekitar 89%, sementara akurasi terendah terjadi pada fold kedua dengan nilai sekitar 86%. Hal ini mengindikasikan bahwa *Naive Bayes* kurang stabil dalam menangani variasi data pada aspek lokasi.

Confussion matrix klasifikasi sentimen pada aspek fasilitas dari model *Naive Bayes* dan *Logistic Regression* ditampilkan pada



Gambar 4.16 *Confusion Matrix* Model *Naive Bayes* pada Klasifikasi Sentimen Aspek Lokasi



Gambar 4.17 *Confusion Matrix Model Logistic Regression* pada Klasifikasi Sentimen Aspek Lokasi

Gambar *confusion matrix* menunjukkan hasil evaluasi model *Naive Bayes* dan *Logistic Regression* menggunakan 5-fold *cross validation*. Pada model *Naive Bayes*, terlihat bahwa prediksi terhadap kelas 1 cukup baik dengan nilai *true positive* yang tinggi pada setiap fold, namun model masih menghasilkan *false negative* dan *false positive* yang relatif lebih besar, sehingga kemampuan dalam mengenali kelas 0 belum optimal. Sementara itu, model *Logistic Regression* menunjukkan performa yang lebih stabil dan akurat. Nilai *true positive* lebih tinggi dibandingkan *Naive Bayes*, dan jumlah *false negative* maupun *false positive* lebih rendah, menandakan model ini lebih efektif dalam membedakan kedua kelas. Secara keseluruhan, *Logistic Regression* memberikan performa yang lebih baik dan konsisten pada semua fold dibandingkan *Naive Bayes*.

Berdasarkan hasil pengujian dan pembahasan dapat disimpulkan bahwa metode ***Logistic Regression*** menunjukkan performa yang lebih baik dibandingkan

Naïve Bayes dalam analisis sentimen berbasis aspek pada pelayanan hotel di Pekanbaru. Hal ini terlihat dari hasil evaluasi menggunakan *Stratified K-Fold Cross Validation*, di mana *Logistic Regression* memperoleh rata-rata akurasi sebesar **91,16%**, sedangkan *Naïve Bayes* mencapai **87,38%**.

Keunggulan *Logistic Regression* ditunjukkan oleh kemampuannya dalam memodelkan hubungan fitur TF-IDF secara lebih optimal, sehingga menghasilkan klasifikasi sentimen yang lebih akurat pada setiap aspek yang dianalisis. Dengan demikian, meskipun kedua metode mampu melakukan klasifikasi sentimen dengan baik, *Logistic Regression* merupakan metode yang paling optimal untuk digunakan pada analisis sentimen berbasis aspek dalam penelitian ini.

4.8 Pengujian Model Klasifikasi

Tahapan ini merupakan tahapan untuk menguji model yang telah di latih menggunakan data baru yang tidak termasuk ke dalam data latih. Tahapan ini menguji 5 data sebagai berikut.

Tabel 4.8 Data Pengujian

No.	Ulasan	Aspek	Sentimen
1.	Resepsionis terlihat bingung saat check-in kami	Pelayanan	Negatif
2.	Gym kecil dan peralatan banyak yang rusak	Fasilitas	Negatif
3.	Kami menerima tagihan yang salah	Pelayanan	Negatif
4.	kamar mandi becek	Fasilitas	Negatif
5.	Staf sangat ramah dan membantu dengan rekomendasi wisata	Pelayanan	Positif

Hasil pengujian model yang telah dilatih pada tahapan sebelumnya menggunakan data pada Tabel 4.8 adalah sebagai berikut.

Tabel 4.9 Hasil Pengujian

No.	Aktual		Prediksi	
	Aspek	Sentimen	Aspek	Sentimen
1	Pelayanan	Negatif	Pelayanan	Negatif
2	Fasilitas	Negatif	Fasilitas	Negatif
3	Pelayanan	Negatif	Pelayanan	Negatif
4	Fasilitas	Negatif	Fasilitas	Negatif
5	Pelayanan	Positif	Pelayanan	Positif

Hasil pengujian menunjukkan bahwa model memiliki kemampuan yang baik dalam mengidentifikasi aspek yang dibahas dalam ulasan hotel serta menentukan polaritas sentimennya. Kinerja yang konsisten antara label aktual dan prediksi pada seluruh data uji mengindikasikan bahwa model telah mampu melakukan generalisasi dengan baik terhadap data baru yang tidak dilibatkan dalam proses pelatihan. Hasil ini menjadi indikator bahwa model siap digunakan pada proses analisis sentimen berbasis aspek.

BAB V

KESIMPULAN

5.1 Kesimpulan

Berdasarkan hasil penelitian yang sudah diuraikan pada bab sebelumnya, maka dapat disimpulkan sebagai berikut.

1. Pemodelan topik menggunakan metode LDA menghasilkan tiga topik dengan nilai *coherence* sebesar 0,51 dan *perplexity* sebesar -6,733. Pemilihan jumlah topik sebanyak tiga didasarkan pada analisis menggunakan *elbow method*. Dari tahapan ini, tiga aspek utama yang berhasil diidentifikasi adalah fasilitas, pelayanan, dan lokasi.
2. *Naïve Bayes* menunjukkan kinerja terbaik dalam klasifikasi aspek dengan akurasi mencapai 88,1%. Sedangkan pada klasifikasi sentimen berbasis aspek, *Logistic Regression* merupakan metode pembelajaran mesin dengan kinerja terbaik dengan akurasi sebesar 93,6% pada aspek fasilitas, 83,3% pada aspek pelayanan, dan 91,2% pada aspek lokasi.
3. Metode *Logistic Regression* terbukti lebih efektif dibandingkan *Naive Bayes* dalam menganalisis sentimen berbasis aspek pada ulasan hotel di Pekanbaru baik dari segi akurasi maupun kestabilan performa model. Sehingga metode *Logistic Regression* lebih direkomendasikan untuk digunakan dalam penelitian analisis sentimen berbasis aspek.

5.2 Saran

1. Penelitian ini hanya menggunakan ulasan dari Google Maps. Maka disarankan untuk memperluas sumber data, contohnya mengambil ulasan dari berbagai platform.
2. Pengembangan metode model penelitian ini hanya membandingkan dua metode yaitu *Naïve Bayes* dan *Logistic Regression*. Disarankan untuk dapat mengeksplorasi metode yang lebih kompleks seperti BERT, LSTM atau IndoBERT.
3. Pada penelitian ini menggunakan metode SMOTE untuk menangani ketidakseimbangan kelas. Lebih baik untuk mencoba pendekatan lain untuk membandingkan hasil dan menentukan metode yang efektif.

DAFTAR PUSTAKA

- Akbar, J., M., T. A., Tolla, Y., Ahmad, A. E., Yaqin, A., & Utami, E. (2023). Pemodelan Topik Menggunakan Latent Dirichlet Allocation pada Ulasan Aplikasi PeduliLindungi. *InComTech : Jurnal Telekomunikasi Dan Komputer*, 13(1). <https://doi.org/10.22441/incomtech.v13i1.15572>
- Amira, S. A., Utama, S., & Fahmi, M. H. (2020). Penerapan Metode Support Vector Machine untuk Analisis Sentimen pada Review Pelanggan Hotel. *Edu Komputika Journal*, 7(2). <https://doi.org/10.15294/edukomputika.v7i2.42608>
- Anandarajan, M., Hill, C., & Nolan, T. (2019a). Term-Document Representation. In M. Anandarajan, C. Hill, & T. Nolan (Eds.), *Practical Text Analytics: Maximizing the Value of Text Data* (pp. 61–73). Springer International Publishing. https://doi.org/10.1007/978-3-319-95663-3_5
- Anandarajan, M., Hill, C., & Nolan, T. (2019b). Text Preprocessing. In M. Anandarajan, C. Hill, & T. Nolan (Eds.), *Practical Text Analytics: Maximizing the Value of Text Data* (pp. 45–59). Springer International Publishing. https://doi.org/10.1007/978-3-319-95663-3_4
- Arianto, D., & Budi, I. (2023). Analisis Sentimen Berbasis Aspek dan Pemodelan Topik pada Candi Borobudur dan Candi Prambanan. *MULTINETICS*, 8(2). <https://doi.org/10.32722/multinetics.v8i2.5056>
- Arora, K., & Rangarajan, A. (2016). *Contrastive Entropy: A new evaluation metric for unnormalized language models*. <http://arxiv.org/abs/1601.00248>
- Assaidi, S. A., & Amin, F. (2022). Analisis Sentimen Evaluasi Pembelajaran Tatap Muka 100 Persen pada Pengguna Twitter menggunakan Metode Logistic Regression. *Jurnal Pendidikan Tambusai*, 6(2).
- Astuti, S. P. (2020). *Analisis Sentimen Berbasis Aspek pada Aplikasi Tokopedia Menggunakan LDA dan Naïve Bayes*. UIN Syarif Hidayatullah.
- Berrar, D. (2018). Cross-validation. In *Encyclopedia of Bioinformatics and Computational Biology: ABC of Bioinformatics* (Vols. 1–3). <https://doi.org/10.1016/B978-0-12-809633-8.20349-X>
- Blei, D., Ng, A., & Jordan, M. (2003). Latent Dirichlet Allocation. *Journal of Machine Learning Research*, 3, 993–1022. <https://doi.org/10.1162/jmlr.2003.3.4-5.993>

- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic minority over-sampling technique. *Journal of Artificial Intelligence Research*, 16. <https://doi.org/10.1613/jair.953>
- Ding, B., Qin, C., Liu, L., Chia, Y. K., Joty, S., Li, B., & Bing, L. (2022). *Is GPT-3 a Good Data Annotator?* <http://arxiv.org/abs/2212.10450>
- Dzisevic, R., & Sesok, D. (2019). Text Classification using Different Feature Extraction Approaches. *2019 Open Conference of Electrical, Electronic and Information Sciences, EStream 2019 - Proceedings*. <https://doi.org/10.1109/eStream.2019.8732167>
- Erfina, A., & Resti Wardani, N. (2022). Analisis Sentimen Perguruan Tinggi Termewah Di Indonesia Menurut Ulasan Google Maps Menggunakan Support Vector Machine (Svm). *Jurnal Manajemen Informatika & Sistem Informasi (MISI)*, 5.
- Fagbohun, O., Harrison, R. M., & Dereventsov, A. (2023). An Empirical Categorization of Prompting Techniques for Large Language Models: A Practitioner's Guide. *Journal of Artificial Intelligence, Machine Learning and Data Science*, 1(4). <https://doi.org/10.51219/jaimld/oluwole-fagbohun/15>
- Glanville, J. (2018). Text mining for information specialists. In J. Craven & P. Levay (Eds.), *Systematic Searching* (pp. 147–170). Facet. <https://doi.org/DOI:10.29085/9781783303755.008>
- Gulati, K., Saravana Kumar, S., Sarath Kumar Boddu, R., Sarvakar, K., Kumar Sharma, D., & Nomani, M. Z. M. (2022). Comparative analysis of machine learning-based classification models using sentiment classification of tweets related to COVID-19 pandemic. *Materials Today: Proceedings*, 51, 38–41. <https://doi.org/https://doi.org/10.1016/j.matpr.2021.04.364>
- Helmayanti, S. A., Hamami, F., & Fa'rifah, R. Y. (2023). PENERAPAN ALGORITMA TF-IDF DAN NAÏVE BAYES UNTUK ANALISIS SENTIMEN BERBASIS ASPEK ULASAN APLIKASI FLIP PADA GOOGLE PLAY STORE. *Jurnal Indonesia : Manajemen Informatika Dan Komunikasi*, 4(3). <https://doi.org/10.35870/jimik.v4i3.415>
- Hickman, L., Thapa, S., Tay, L., Cao, M., & Srinivasan, P. (2022). Text Preprocessing for Text Mining in Organizational Research: Review and Recommendations. *Organizational Research Methods*, 25(1). <https://doi.org/10.1177/1094428120971683>
- Ibrohim, M. O., & Budi, I. (2019). *Multi-label Hate Speech and Abusive Language Detection in Indonesian Twitter*. <https://doi.org/10.18653/v1/w19-3506>

- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). Resampling Methods. In G. James, D. Witten, T. Hastie, & R. Tibshirani (Eds.), *An Introduction to Statistical Learning: with Applications in R* (pp. 197–223). Springer US. https://doi.org/10.1007/978-1-0716-1418-1_5
- Kapadia, S. (2019). *Evaluate Topic Models : Latent Dirichlet Allocation (LDA)*. <https://Towardsdatascience.Com/Evaluate-Topic-Model-in-Python-Latent-Dirichlet-Allocation-Lda-7d57484bb5d0>.
- Kherwa, P., & Bansal, P. (2020). Topic Modeling: A Comprehensive Review. *EAI Endorsed Transactions on Scalable Information Systems*, 7(24). <https://doi.org/10.4108/eai.13-7-2018.159623>
- Kochmar, E. (2022). *Getting Started with Natural Language Processing*. Manning.
- Korencic, D., Ristov, S., Repar, J., & Snajder, J. (2021). A Topic Coverage Approach to Evaluation of Topic Models. *IEEE Access*, 9. <https://doi.org/10.1109/ACCESS.2021.3109425>
- Lamba, M., & Madhusudhan, M. (2022). Text Mining for Information Professionals: An Uncharted Territory. In *Text Mining for Information Professionals: An Uncharted Territory*. <https://doi.org/10.1007/978-3-030-85085-2>
- Li, J., Tang, T., Zhao, W. X., Nie, J.-Y., & Wen, J.-R. (2022). *Pre-trained Language Models for Text Generation: A Survey*.
- Li, Y. (2023). A Practical Survey on Zero-shot Prompt Design for In-context Learning. *International Conference Recent Advances in Natural Language Processing, RANLP*. https://doi.org/10.26615/978-954-452-092-2_069
- Liang, H., Sun, X., Sun, Y., & Gao, Y. (2017). Text feature extraction based on deep learning: a review. In *Eurasip Journal on Wireless Communications and Networking* (Vol. 2017, Issue 1). <https://doi.org/10.1186/s13638-017-0993-1>
- Liu, B. (2020). Sentiment Analysis: Mining Opinions, Sentiments, and Emotions. In *Studies in Natural Language Processing* (2nd ed.). Cambridge University Press. <https://doi.org/DOI: 10.1017/9781108639286>
- Muralidhar, K. (2021). *The right way of using SMOTE with Cross-validation*. <https://Towardsdatascience.Com/the-Right-Way-of-Using-Smote-with-Cross-Validation-92a8d09d00c7>.
- Narkhede, S. (2018). *Understanding Confusion Matrix*. <https://Towardsdatascience.Com/Understanding-Confusion-Matrix-A9ad42dcfd62>.

- Nasution, A. H., & Onan, A. (2024). ChatGPT Label: Comparing the Quality of Human-Generated and LLM-Generated Annotations in Low-Resource Language NLP Tasks. *IEEE Access*, 12, 71876–71900. <https://doi.org/10.1109/ACCESS.2024.3402809>
- Parasati, W., Abdurrachman Bachtiar, F., & Setiawan, N. Y. (2020). Analisis Sentimen Berbasis Aspek pada Ulasan Pelanggan Restoran Bakso President Malang dengan Metode Naïve Bayes Classifier. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 4(4).
- Pratama, Y., Murdiansyah, D. T., & Lhaksana, K. M. (2023). Analisis Sentimen Kendaraan Listrik Pada Media Sosial Twitter Menggunakan Algoritma Logistic Regression dan Principal Component Analysis. *JURNAL MEDIA INFORMATIKA BUDIDARMA*, 7(1). <https://doi.org/10.30865/mib.v7i1.5575>
- Rahmawati, I., Rika Fitriani, T., No'eman, A., & Yusuf, A. Y. P. (2023). Analisis Sentimen Menggunakan Algoritma Logistic Regression Pada Penerbangan Lion Air berdasarkan Ulasan Platform Online. *Jurnal Riset Informatika Dan Teknologi Informasi*, 1(1). <https://doi.org/10.58776/jriti.v1i1.60>
- Raschka, S. (2018). *Model Evaluation, Model Selection, and Algorithm Selection in Machine Learning*.
- Ray, S. (2019). A Quick Review of Machine Learning Algorithms. *Proceedings of the International Conference on Machine Learning, Big Data, Cloud and Parallel Computing: Trends, Perspectives and Prospects, COMITCon 2019*. <https://doi.org/10.1109/COMITCon.2019.8862451>
- Reynolds, L., & McDonell, K. (2021). Prompt Programming for Large Language Models: Beyond the Few-Shot Paradigm. *Conference on Human Factors in Computing Systems - Proceedings*. <https://doi.org/10.1145/3411763.3451760>
- Röder, M., Both, A., & Hinneburg, A. (2015). Exploring the space of topic coherence measures. *WSDM 2015 - Proceedings of the 8th ACM International Conference on Web Search and Data Mining*. <https://doi.org/10.1145/2684822.2685324>
- Schulhoff, S., Ilie, M., Balepur, N., Kahadze, K., Liu, A., Si, C., Li, Y., Gupta, A., Han, H., Schulhoff, S., Dulepet, P. S., Vidyadhara, S., Ki, D., Agrawal, S., Pham, C., Kroiz, G., Li, F., Tao, H., Srivastava, A., ... Resnik, P. (2025). *The Prompt Report A Systematic Survey of Prompt Engineering Techniques*. <https://arxiv.org/abs/2406.06608>
- Setiawan, R. A., Estetikha, A. K. A., Nurharyanto, E. M. O., Asmara, Y., & Wahyudi, A. (2022). Analisis Sentimen Hotel di Nusa Tenggara Barat Menggunakan Algoritma SVM. *Jurnal Informasi Interaktif*, 7(1).

- Setiowati, Y., & Helen, A. (2018). KLASIFIKASI ANALISIS SENTIMEN MENGENAI HOTEL DI YOGYAKARTA. *SCAN - Jurnal Teknologi Informasi Dan Komunikasi*, 13(1). <https://doi.org/10.33005/scan.v13i1.1052>
- Suciati, A., & Budi, I. (2019). Aspect-based Opinion Mining for Code-Mixed Restaurant Reviews in Indonesia. *Proceedings of the 2019 International Conference on Asian Language Processing, IALP 2019*. <https://doi.org/10.1109/IALP48816.2019.9037689>
- Teng, Z., & Zhang, Y. (2021). *Natural Language Processing: A Machine Learning Perspective*. Cambridge University Press. <https://doi.org/DOI:10.1017/9781108332873>
- Tijare, P., & Rani, P. J. (2020). Exploring popular topic models. *Journal of Physics: Conference Series*, 1706(1). <https://doi.org/10.1088/1742-6596/1706/1/012171>
- Tran Xuan, T., Ba, H., & Huynh, V.-N. (2019). *Measuring Hotel Review Sentiment: An Aspect-Based Sentiment Analysis Approach* (pp. 393–405). https://doi.org/10.1007/978-3-030-14815-7_33
- Tripathy, A., Agrawal, A., & Rath, S. K. (2016). Classification of sentiment reviews using n-gram machine learning approach. *Expert Systems with Applications*, 57. <https://doi.org/10.1016/j.eswa.2016.03.028>
- Vujović, Ž. (2021). Classification Model Evaluation Metrics. *International Journal of Advanced Computer Science and Applications*, 12(6). <https://doi.org/10.14569/IJACSA.2021.0120670>
- Wongvorachan, T., He, S., & Bulut, O. (2023). A Comparison of Undersampling, Oversampling, and SMOTE Methods for Dealing with Imbalanced Classification in Educational Data Mining. *Information (Switzerland)*, 14(1). <https://doi.org/10.3390/info14010054>
- Yang, F. J. (2018). An implementation of naive bayes classifier. *Proceedings - 2018 International Conference on Computational Science and Computational Intelligence, CSCI 2018*. <https://doi.org/10.1109/CSCI46756.2018.00065>
- Yenduri, G., M, R., G, C. S., Y, S., Srivastava, G., Maddikunta, P. K. R., G, D. R., Jhaveri, R. H., B, P., Wang, W., Vasilakos, A. V., & Gadekallu, T. R. (2023). *Generative Pre-trained Transformer: A Comprehensive Review on Enabling Technologies, Potential Applications, Emerging Challenges, and Future Directions*. <http://arxiv.org/abs/2305.10435>
- Zhihao, P., Fenglong, Y., & Xucheng, L. (2019). Comparison of the Different Sampling Techniques for Imbalanced Classification Problems in Machine Learning. *2019 11th International Conference on Measuring Technology and*

Mechatronics Automation (ICMTMA), 431–434.
<https://doi.org/10.1109/ICMTMA.2019.00101>

UNIVERSITAS ISLAM RIAU



SURAT KEPUTUSAN DEKAN FAKULTAS TEKNIK UNIVERSITAS ISLAM RIAU
NOMOR : 0713/KPTS/FT-UIR/2025
TENTANG PENGANGKATAN TIM PEMBIMBING PENELITIAN DAN PENYUSUNAN SKRIPSI

DEKAN FAKULTAS TEKNIK

- Membaca : Surat Ketua Program Studi Teknik Informatika Nomor : 263/TA-TI/FT/2024 tentang persetujuan dan usulan pengangkatan Tim Pembimbing penelitian dan penyusunan Skripsi.
- Menimbang : 1. Bahwa untuk menyelesaikan perkuliahan bagi mahasiswa Fakultas Teknik perlu membuat Skripsi.
2. Untuk itu perlu ditunjuk Tim Pembimbing penelitian dan penyusunan Skripsi yang diangkat dengan Surat Keputusan Dekan.
- Mengingat : 1. Undang - Undang Nomor 12 Tahun 2012 Tentang Pendidikan Tinggi
2. Peraturan Presiden Republik Indonesia Nomor 8 Tahun 2012 Tentang Kerangka Kualifikasi Nasional Indonesia
3. Peraturan Pemerintah Republik Indonesia Nomor 37 Tahun 2009 Tentang Dosen
4. Peraturan Pemerintah Republik Indonesia Nomor 66 Tahun 2010 Tentang Pengelolaan dan Penyelenggaraan Pendidikan
5. Peraturan Menteri Pendidikan Nasional Nomor 63 Tahun 2009 Tentang Sistem Penjaminan Mutu Pendidikan
6. Peraturan Menteri Pendidikan dan Kebudayaan Republik Indonesia Nomor 49 Tahun 2014 Tentang Standar Nasional Pendidikan Tinggi
7. Statuta Universitas Islam Riau Tahun 2018
8. Peraturan Universitas Islam Riau Nomor 001 Tahun 2018 Tentang Ketentuan Akademik Bidang Pendidikan Universitas Islam Riau

MEMUTUSKAN

- Menetapkan : 1. Mengangkat saudara-saudara yang namanya tersebut dibawah ini sebagai Tim Pembimbing Penelitian & penyusunan Skripsi Mahasiswa Fak. Teknik Program Studi Teknik Informatika.

No	Nama	Pangkat	Jabatan
1.	Anggi Hanafiah, S.Kom., M.Kom.	Lektor	Pembimbing

2. Mahasiswa yang akan dibimbing :

Nama : Fadhia Haya
NPM : 203510487
Program Studi : Teknik Informatika
Jenjang Pendidikan : Strata Satu (S1)
Judul Skripsi : Analisis Berbasis Aspek Pada Pelayanan Hotel Di Pekanbaru
Menggunakan Metode Naive Bayes Dan Logistic Regression

3. Keputusan ini mulai berlaku pada tanggal ditetapkannya dengan ketentuan bila terdapat kekeliruan dikemudian hari segera ditinjau kembali.

Ditetapkan di : Pekanbaru
Pada Tanggal : 20 Jumadil Awal 1447 H
11 November 2025 M

Dekan,



Dr. Ir. KURNIA HASTUTI, S.T, M.T

NPK : 1008057102

Tembusan disampaikan :

1. Yth. Bapak Rektor UIR di Pekanbaru.
2. Yth. Sdr. Ketua Program Studi Teknik Informatika FT-UIR
3. Arsip



YAYASAN LEMBAGA PENDIDIKAN ISLAM (YLPI) RIAU
UNIVERSITAS ISLAM RIAU

F.A.3.10

Jalan Kahrudin Nasution No. 113 P. Marpoan Pekanbaru Riau Indonesia – Kode Pos: 28284
Telp. +62 761 674674 Fax. +62 761 674834 Website: www.uir.ac.id Email: info@uir.ac.id

KARTU BIMBINGAN TUGAS AKHIR
SEMESTER GANJIL TA 2024/2025

NPM : 203510487
Nama Mahasiswa : FADHIA HAYA
Dosen Pembimbing : 1. ANGGI HANAFIAH S.Kom., M.Kom 2.
Program Studi : TEKNIK INFORMATIKA
Judul Tugas Akhir : Analisis Sentimen Berbasis Aspek pada Pelayanan Hotel di Pekanbaru
Menggunakan Metode Naive Bayes dan Logistic Regression
Judul Tugas Akhir (Bahasa Inggris) : Aspect-based Sentiment analysis on hotel services in Pekanbaru using Naive Bayes and Logistic Regression methods
Lembar Ke : 1

NO	Hari/Tanggal Bimbingan	Materi Bimbingan	Hasil / Saran Bimbingan	Paraf Dosen Pembimbing
1	Selasa 11 juni 2024	Judul skripsi dan metode	Acc judul dan metode, cari data set	
2	Kamis 22 agustus 2024	Pengumpulan data	Acc data	
3	Selasa 17 sep 2024	Bab 1	Revisi bab 1	
4	Kamis 17 okt 2024	Bab 1 dan Bab 2	Revisi bab 2 dan lanjut bab 3	
5	Selasa 29 okt 2024	Bab 3	Revisi alur penelitian	
6	Kamis 16 jan 2025	Bab 3	Perbaiki penulisan	
7	Selasa 18 feb 2025	Bab 3	Acc laporan	
8	Kamis 17 april 2025	Hasil akhir dari kodingan analisis sentimen	Acc	

Pekanbaru, 22 April 2025
Wakil Dekan I/Ketua Departemen/Ketua Prodi



MJAZNTEWNDG3



Catatan :

- Lama bimbingan Tugas Akhir/ Skripsi maksimal 2 semester sejak TMT SK Pembimbing diterbitkan
- Kartu ini harus dibawa setiap kali berkonsultasi dengan pembimbing dan HARUS dicetak kembali setiap memasuki semester baru melalui SIKAD
- Saran dan koreksi dari pembimbing harus ditulis dan diparaf oleh pembimbing
- Setelah skripsi disetujui (ACC) oleh pembimbing, kartu ini harus ditandatangani oleh Wakil Dekan I/ Kepala departemen/Ketua prodi
- Kartu kendali bimbingan asli yang telah ditandatangani diserahkan kepada Ketua Program Studi dan kopiannya dilampirkan pada skripsi.
- Jika jumlah pertemuan pada kartu bimbingan tidak cukup dalam satu halaman, kartu bimbingan ini dapat di download kembali melalui SIKAD



YAYASAN LEMBAGA PENDIDIKAN ISLAM (YLPI) RIAU
UNIVERSITAS ISLAM RIAU

F.A.3.10

Jalan Kaharuddin Nasution No. 113 P. Marpoan Pekanbaru Riau Indonesia – Kode Pos: 28284
Telp. +62 761 674674 Fax. +62 761 674834 Website: www.uir.ac.id Email: info@uir.ac.id

KARTU BIMBINGAN TUGAS AKHIR
SEMESTER GANJIL TA 2024/2025

NPM : 203510487
Nama Mahasiswa : FADHIA HAYA
Dosen Pembimbing : 1. ANGGI HANAFIAH S.Kom., M.Kom 2.
Program Studi : TEKNIK INFORMATIKA
Judul Tugas Akhir : Analisis Sentimen Berbasis Aspek pada Pelayanan Hotel di Pekanbaru
Menggunakan Metode Naive Bayes dan Logistic Regression
Judul Tugas Akhir
(Bahasa Inggris) : Aspect-based Sentiment analysis on hotel services in Pekanbaru using Naive
Bayes and Logistic Regression methods
Lembar Ke : 1

NO	Hari/Tanggal Bimbingan	Materi Bimbingan	Hasil / Saran Bimbingan	Paraf Dosen Pembimbing
1	Rabu 19 Nov 2025	Bimbingan bab 4 dan 5	Penambahan penjelasan gambar pada bab 4	
2	Jumat 5 Des 2025	Penambahan penjelasan bab 4	Menambahkan pegujian data pada bab 4	
3	Jumat 12 Des 2025	Penambahan pengujian data	Acc	

Pekanbaru, 15 November 2025
Wakil Dekan I/Ketua Departemen/Ketua Prodi



MJAZNTEWNDG3



Catatan :

- Lama bimbingan Tugas Akhir/ Skripsi maksimal 2 semester sejak TMT SK Pembimbing diterbitkan
- Kartu ini harus dibawa setiap kali berkonsultasi dengan pembimbing dan HARUS dicetak kembali setiap memasuki semester baru melalui SIKAD
- Saran dan koreksi dari pembimbing harus ditulis dan diparaf oleh pembimbing
- Setelah skripsi disetujui (ACC) oleh pembimbing, kartu ini harus ditandatangani oleh Wakil Dekan I/ Kepala departemen/Ketua prodi
- Kartu kendali bimbingan asli yang telah ditandatangani diserahkan kepada Ketua Program Studi dan kopiannya dilampirkan pada skripsi.
- Jika jumlah pertemuan pada kartu bimbingan tidak cukup dalam satu halaman, kartu bimbingan ini dapat di download kembali melalui SIKAD

SURAT KEPUTUSAN DEKAN FAKULTAS TEKNIK UNIVERSITAS ISLAM RIAU
NOMOR : 0028/KPTS/FT-UIR/2026
TENTANG PENETAPAN DOSEN PENGUJI SKRIPSI MAHASISWA FAK. TEKNIK UNIV. ISLAM RIAU

DEKAN FAKULTAS TEKNIK

- Menimbang** : 1. Bahwa untuk menyelesaikan studi S.1 bagi mahasiswa Fakultas Teknik Univ. Islam Riau dilaksanakan Ujian Skripsi/Komprehensif sebagai tugas akhir. Untuk itu perlu ditetapkan mahasiswa yang telah memenuhi syarat untuk ujian dimaksud serta dosen penguji.
2. Bahwa penetapan mahasiswa yang memenuhi syarat dan dosen penguji yang bersangkutan perlu ditetapkan dengan Surat Keputusan Dekan.
- Mengingat** : 1. Undang - Undang Nomor 12 Tahun 2012 Tentang Pendidikan Tinggi
2. Peraturan Presiden Republik Indonesia Nomor 8 Tahun 2012 Tentang Kerangka Kualifikasi Nasional Indonesia
3. Peraturan Pemerintah Republik Indonesia Nomor 37 Tahun 2009 Tentang Dosen
4. Peraturan Pemerintah Republik Indonesia Nomor 66 Tahun 2010 Tentang Pengelolaan dan Penyelenggaraan Pendidikan
5. Peraturan Menteri Pendidikan Nasional Nomor 63 Tahun 2009 Tentang Sistem Penjaminan Mutu Pendidikan
6. Peraturan Menteri Pendidikan dan Kebudayaan Republik Indonesia Nomor 49 Tahun 2014 Tentang Standar Nasional Pendidikan Tinggi
7. Statuta Universitas Islam Riau Tahun 2018
8. Peraturan Universitas Islam Riau Nomor 001 Tahun 2018 Tentang Ketentuan Akademik Bidang Pendidikan Universitas Islam Riau

MEMUTUSKAN

- Menetapkan** : 1. Mahasiswa Fakultas Teknik Universitas Islam Riau yang tersebut namanya dibawah ini :
- | | |
|--------------------|---|
| Nama | : Fadhia Haya |
| NPM | : 203510487 |
| Program Studi | : Teknik Informatika |
| Jenjang Pendidikan | : Strata Satu (S1) |
| Judul Skripsi | : Analisis Sentimen Berbasis Aspek pada Pelayanan Hotel di Pekanbaru Menggunakan Metode Naïve Bayes dan Logistic Regression |
2. Penguji Skripsi/Komprehensif mahasiswa tersebut terdiri dari :
- | | |
|---|-----------------------------------|
| 1. Anggi Hanafiah, S.Kom., M.Kom | Sebagai Ketua Merangkap Penguji |
| 2. Dr. Arbi Haza Nasution, B.IT., M.IT. | Sebagai Anggota Merangkap Penguji |
| 3. Rizky Wandri, S.Kom., M.Kom | Sebagai Anggota Merangkap Penguji |
3. Laporan hasil ujian serta berita acara telah sampai kepada Pimpinan Fakultas selambat-lambatnya 1(satu) bulan setelah ujian dilaksanakan.
4. Keputusan ini mulai berlaku pada tanggal ditetapkannya dengan ketentuan bila terdapat kekeliruan dikemudian hari segera ditinjau kembali.

KUTIPAN : Disampaikan kepada yang bersangkutan untuk dapat dilaksanakan dengan sebaik-baiknya.

Ditetapkan di : Pekanbaru
Pada Tanggal : 25 Rajab 1447 H
14 Januari 2026 M

Dekan,



Dr. Ir. KURNIA HASTUTI, S.T, M.T

NPK : 1008057102



YAYASAN LEMBAGA PENDIDIKAN ISLAM (YLPI) RIAU
UNIVERSITAS ISLAM RIAU
FAKULTAS TEKNIK
PROGRAM STUDI TEKNIK INFORMATIKA

Jalan Kaharuddin Nasution No. 113 P. Marpoyan Pekanbaru Riau Indonesia – Kode Pos: 28284
Telp. +62 761 674674 Website: www.eng.uir.ac.id Email: fakultas_teknik@uir.ac.id

BERITA ACARA UJIAN SKRIPSI

Berdasarkan Surat Keputusan Dekan Fakultas Teknik Universitas Islam Riau, Pekanbaru, tanggal 14 Januari 2026, Nomor: 0028/KPTS/FT-UIR/2026, maka pada hari Kamis tanggal 15 Januari 2026, telah dilaksanakan Ujian Skripsi Program Studi Teknik Informatika Fakultas Teknik Universitas Islam Riau, Jenjang Studi S1, Tahun Akademik 2025/2026 berikut ini.

1. Nama : Fadhia Haya
2. NPM : 203510487
3. Judul Skripsi : Analisis Sentimen Berbasis Aspek pada Pelayanan Hotel di Pekanbaru Menggunakan Metode Naïve Bayes dan Logistic Regression
4. Waktu Ujian : 10.00 WIB s.d. Selesai
5. Tempat Pelaksanaan Ujian : Ruang Sidang Fakultas Teknik UIR

Dengan keputusan Hasil Ujian Skripsi:

~~Lulus~~*/ Lulus dengan Perbaikan*/ ~~Tidak Lulus~~*

**Coret yang tidak perlu.*

Nilai Ujian:

Nilai Ujian Angka = 71,95 Nilai Huruf = (B+)

Tim Penguji Skripsi.

No	Nama	Jabatan	Tanda Tangan
1	Anggi Hanafiah, S.Kom., M.Kom	Ketua	1.
2	Dr. Arbi Haza Nasution, B.IT., M.IT.	Anggota	2.
3	Rizky Wandri, S.Kom., M.Kom	Anggota	3.

Panitia Ujian
Ketua,

Anggi Hanafiah, S.Kom., M.Kom.
NIDN. 1014028904

Pekanbaru, 15 Januari 2026

Mengetahui,
Dekan Fakultas Teknik



Dr. Ir. Kurnia Hastuti, S.T., M.T.
NIDN. 1008057102



UNIVERSITAS ISLAM RIAU

FAKULTAS TEKNIK

الْجَامِعَةُ الْإِسْلَامِيَّةُ الرَّيْوِيَّةُ

Alamat: Jalan Kaharuddin Nasution No.113, Marpoyan, Pekanbaru, Riau, Indonesia - 28284
Telp. +62 761 674674 Email: fakultas_teknik@uir.ac.id Website: www.eng.uir.ac.id

SURAT KETERANGAN BEBAS PLAGIAT

Nomor: 653/A-UIR/5-T/2025

Fakultas Teknik Universitas Islam Riau menerangkan bahwa Mahasiswa/i dengan identitas berikut:

Nama : **FADHIA HAYA**
NPM : 203510487
Program Studi : Teknik Informatika
Jenjang Pendidikan : Strata Satu (S1)
Judul Skripsi TA : ANALISIS SENTIMEN BERBASIS ASPEK PADA PELAYANAN HOTEL DI PEKANBARU MENGGUNAKAN METODE NAIVE BAYES DAN LOGISTIC REGRESSION

Dinyatakan **Bebas Plagiat**, berdasarkan hasil pengecekan pada Turnitin menunjukkan angka **Similarity Index < 30%** sesuai dengan peraturan Universitas Islam Riau yang berlaku.

Demikian surat keterangan ini dibuat untuk dapat dipergunakan sebagaimana mestinya.

Mengetahui,

Kaprodi. Teknik Informatika

Assoc. Prof. Dr. Ir. Hj. Des Suryani, M.Sc.

Pekanbaru, 30 December 2025 M

10 Rojab 1447 H

Staff Pemeriksa

Syatra Safira Anjani, S.Pd