

**ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN WISATA
JATIM PARK GROUP MENGGUNAKAN *RANDOM FOREST***

SKRIPSI

Oleh:
MOCH HASBI ASHIDQY
NIM. 21605110104



**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

**ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN WISATA
JATIM PARK GROUP MENGGUNAKAN *RANDOM FOREST***

SKRIPSI

Diajukan kepada :
Universitas Islam Negeri Maulana Malik Ibrahim Malang
Untuk memenuhi Salah Satu Persyaratan dalam
Memperoleh Gelar Sarjana Komputer (S.Kom)

Oleh:
MOCH HASBI ASHIDQY
NIM. 210605110104

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

HALAMAN PERSETUJUAN

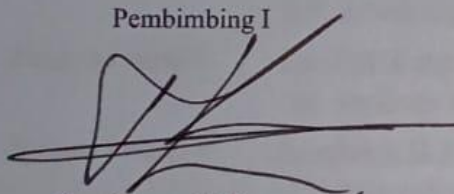
ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN WISATA
JATIM PARK GROUP MENGGUNAKAN RANDOM FOREST

SKRIPSI

Oleh :
MOCH HASBI ASHIDQY
NIM. 21605110104

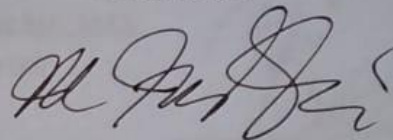
Telah Diperiksa dan Disetujui untuk Diuji :
Tanggal : 02 Juni 2026

Pembimbing I



Supriyono, M. Kom
NIP. 19841010 201903 1 012

Pembimbing 2



Dr. H. Mochamad Imamudin, Lc., M.A
NIP. 19740602 200901 1 010

Mengetahui

Ketua Program Studi Teknik Informatika

Fakultas Sains dan Teknologi

Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyono, M. Kom
NIP. 19841010 201903 1 012

HALAMAN PENGESAHAN

ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN WISATA JATIM PARK GROUP MENGGUNAKAN RANDOM FOREST

SKRIPSI

Oleh :
MOCH HASBI ASHIDQY
NIM. 21605110104

Telah Dipertahankan di Depan Dewan Penguji Skripsi
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan
Untuk Memperoleh Gelar Sarjana Komputer (S.Kom)
Tanggal : 08 Juni 2026

Susunan Dewan Penguji

Ketua Penguji	: <u>Prof. Dr. Ir. Muhammad Faisal, M.T</u> NIP. 19740510 200501 1 007
Anggota Penguji I	: <u>Nur Fitriyah Ayu Tunjung Sari, M.Cs</u> NIP. 19911226 202012 2 001
Anggota Penguji II	: <u>Supriyono, M. Kom</u> NIP. 19841010 201903 1 012
Anggota Penguji III	: <u>Dr. H. Mochamad Imamudin, Lc., M.A</u> NIP. 19740602 200901 1 010

()
()
()
()
()
()

Mengetahui dan Mengesahkan,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyono, M. Kom
NIP. 19841010 201903 1 012

PERNYATAAN KEASLIAN TULISAN

Saya yang bertanda tangan di bawah ini :

Nama : Moch. Hasbi Ashidqy
NIM : 210605110104
Fakultas/ Program Studi : Sains dan Teknologi / Teknik Informatika
Judul Skripsi : Analisis Sentimen Berbasis Aspek Pada Ulasan Wisata Jatim Park Group Menggunakan *Random Forest*

Menyatakan dengan sebenarnya bahwa Skripsi yang saya tulis ini benar-benar merupakan hasil karya saya sendiri, bukan merupakan pengambil alihan data, tulisan, atau pikiran orang lain yang saya akui sebagai hasil tulisan atau pikiran saya sendiri, kecuali dengan mencantumkan sumber cuplikan pada daftar Pustaka

Apabila dikemudian hari terbukti atau dapat dibuktikan skripsi ini merupakan hasil jiplakan, maka saya bersedia menerima sanksi atas perbuatan tersebut.

Malang, 1 Juni 2026
Yang membuat pernyataan,



Moch. Hasbi Ashidqy
NIM.210605110104

MOTTO

*"Barang siapa belum pernah merasakan pahitnya menuntut ilmu walau sesaat,
maka ia akan menelan hinanya kebodohan sepanjang hidupnya."*

(Imam Syafi'i)

"Skripsi yang baik adalah skripsi yang selesai"

*"Semua ahli awalnya adalah seorang pemula. Yang membedakan mereka
hanyalah konsistensi untuk terus melangkah, kerendahan hati untuk terus belajar,
dan keberanian untuk tidak menyerah pada kegagalan."*

HALAMAN PERSEMBAHAN

Sembah sujud dan syukur sedalam-dalamnya saya panjatkan kehadiran Allah *Subhanahu wa Ta'ala* atas segala limpahan rahmat serta hamba-Mu ini berhasil menyelesaikan perjalanan panjang yang penuh tantangan ini. Karya sederhana ini saya persembahkan setulus hati kepada kedua orang tua tercinta, yang doa-doanya menjadi kekuatan terbesar dan kasih sayangnya menjadi alasan utama saya untuk tidak pernah menyerah meskipun rintangan silih berganti datang menghadang. Keberhasilan ini bukanlah milik saya pribadi, melainkan buah dari setiap tetes keringat, pengorbanan, dan kesabaran tiada henti yang telah mereka berikan sepanjang hidup demi melihat senyum bahagia di wajah saya hari ini.

Terima kasih pula saya haturkan kepada seluruh keluarga besar, sahabat seperjuangan, serta bapak dan ibu dosen yang telah memberikan bimbingan, dukungan moral, maupun kritik membangun selama proses penyusunan skripsi ini berlangsung. Kalian adalah saksi hidup atas segala tawa, air mata, serta kerja keras yang tertuang dalam setiap lembar tulisan ini hingga mencapai titik akhir. Semoga karya ini dapat menjadi langkah awal yang baik untuk masa depan yang lebih cerah, serta mampu memberikan manfaat bagi banyak orang sebagaimana harapan yang selalu kita semogakan bersama-sama dalam setiap bait doa.

KATA PENGANTAR

Bismillahirrahmaanirrahiim, Assalamu'alaikum Warahmatullahi Wabarakatuh.

Segala puji dan syukur yang tak terhingga penulis panjatkan kepada Allah *Subhanahu wa Ta'ala* atas rahmat, hidayah, dan petunjuk-Nya. Berkat karunia-Nya, penulis dapat menyelesaikan penyusunan skripsi yang berjudul "Analisis Sentimen Berbasis Aspek Pada Ulasan Wisata Jatim Park Group Menggunakan Random Forest" sebagai salah satu syarat untuk memperoleh gelar Sarjana Komputer pada Program Studi Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.

Shalawat serta salam senantiasa tercurahkan kepada junjungan kita Nabi Muhammad, yang telah membawa umat manusia dari zaman kegelapan menuju zaman yang penuh dengan ilmu pengetahuan seperti saat ini.

Penulis menyadari bahwa penyusunan skripsi ini tidak terlepas dari bantuan, bimbingan, dan dukungan berbagai pihak. Oleh karena itu, dengan penuh rendah hati, penulis mengucapkan terima kasih yang sebesar-besarnya kepada:

1. Prof. Dr. Hj. Ilfi Nur Diana, M.Si., CAHRM., CRMP, selaku rektor Universitas Islam Negeri Maulana Malik Ibrahim Malang.
2. Dr. Agus Mulyono, M.Kes., selaku dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.
3. Supriyono, M. Kom, selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang sekaligus dosen pembimbing I, atas motivasi dan arahan yang diberikan selama proses perkuliahan.

4. Dr. Totok Chamidy, M. Kom, selaku Wali Dosen, atas bimbingan, arahan, dan perhatian yang diberikan kepada penulis selama masa perkuliahan.
5. Dr. H. Mochamad Imamudin, Lc., M.A, selaku dosen pembimbing II, yang dengan penuh kesabaran membimbing penulis hingga selesai.
6. Prof. Dr. Ir. Muhammad Faisal, M.T, dan Nur Fitriyah Ayu Tunjung Sari, M.Cs, selaku dosen penguji saya, yang telah memberikan kritik dan saran yang berharga selama proses ujian skripsi.
7. Kedua orang tua penulis yang tercinta dan tersayang, Ayahanda dan Ibunda atas segala doa, pengorbanan, kasih sayang, dan dukungan moral maupun materiil yang tidak pernah putus diberikan kepada penulis. Semoga Allah *Subhanahu wa Ta'ala* senantiasa melimpahkan kesehatan, kebahagiaan, dan keberkahan kepada keduanya.
8. Saudara-saudara, atas segala doa dan dukungan yang senantiasa diberikan kepada penulis selama menempuh masa studi.
9. Seluruh teman-teman Program Studi Teknik Informatika angkatan 2021, atas kebersamaan, dukungan, dan semangat yang selalu diberikan kepada penulis selama menempuh perkuliahan hingga penyelesaian skripsi ini.
10. Semua pihak yang tidak dapat penulis sebutkan satu per satu, yang telah memberikan bantuan dan dukungan dalam penyelesaian skripsi ini. Semoga Allah *Subhanahu wa Ta'ala* membalas kebaikan kalian semua.

Penulis menyadari bahwa skripsi ini masih jauh dari sempurna. Oleh karena itu, kritik dan saran yang membangun dari berbagai pihak sangat penulis harapkan demi perbaikan di masa mendatang. Semoga skripsi ini dapat memberikan manfaat

bagi penulis khususnya, serta bagi pembaca dan perkembangan ilmu pengetahuan pada umumnya.

Wassalamu'alaikum Warahmatullahi Wabarakatuh

Malang, 1 Juni 2026

Penulis

DAFTAR ISI

HALAMAN PERSETUJUAN	iii
HALAMAN PENGESAHAN.....	iii
PERNYATAAN KEASLIAN TULISAN.....	v
MOTTO	vi
HALAMAN PERSEMBAHAN	vii
KATA PENGANTAR.....	viii
DAFTAR ISI	xi
DAFTAR TABEL	xiv
DAFTAR GAMBAR.....	xv
ABSTRAK	xvi
ABSTRACT.....	xvii
المُنْخَص	xviii
BAB I PENDAHULUAN.....	1
1.1 Latar belakang.....	1
1.2 Rumusan Masalah	5
1.3 Tujuan Penelitian	6
1.4 Batasan Penelitian	6
1.5 Manfaat Penelitian	7
BAB II STUDI PUSTAKA	8
2.1 Konsep 3A dalam Pariwisata	8
2.2 Analisis Sentimen Berbasis Aspek.....	12
2.3 <i>Preprocessing Data</i>	15
2.3.1 <i>Cleansing</i>	15
2.3.2 <i>Case Folding</i>	16
2.3.3 <i>Normalization</i>	16
2.3.4 <i>Tokenization</i>	17
2.3.5 <i>Stopword Removal</i>	17
2.3.6 <i>Stemming</i>	17
2.4 <i>Term Frequency–Inverse Document Frequency (TF-IDF)</i>	18
2.4.1 <i>Term Frequency (TF)</i>	19
2.4.2 <i>Inverse Document Frequency (IDF)</i>	19
2.4.3 Perhitungan Bobot TF-IDF	20

2.5	<i>Random Forest</i>	21
2.6	Evaluasi Performa Klasifikasi.....	25
2.6.1	<i>Confusion Matrix</i>	25
2.6.2	Metrik Pengukuran Kinerja.....	26
2.7	Penelitian Terkait	28
BAB III	DESAIN PENELITIAN.....	32
3.1	Jenis Penelitian.....	32
3.2	Perancangan Alur Penelitian	33
3.3	Pengumpulan Data	34
3.3.1	Teknik Pengumpulan Data	34
3.3.2	Karakteristik dan Riwayat Dataset.....	35
3.3.3	Anotasi Data.....	37
3.4	<i>Preprocessing</i>	39
3.5	Ekstraksi Fitur TF-IDF.....	45
3.6	Algoritma <i>Random Forest</i>	51
3.6.1	<i>Bootstrap Sampling</i>	51
3.6.2	<i>Random Feature Selection</i>	53
3.6.3	Pembentukan Pohon Keputusan.....	53
3.6.4	Mekanisme Klasifikasi (Majority Voting).....	54
3.7	Evaluasi Pengujian.....	56
3.8	Skenario Pengujian.....	59
BAB IV	HASIL DAN PEMBAHASAN	63
4.1	Hasil Uji Coba.....	63
4.1.1	Hasil Uji Skenario 1	64
4.1.2	Hasil Uji Skenario 2	68
4.1.3	Hasil Uji Skenario 3	71
4.1.4	Hasil Uji Skenario 4.....	74
4.2	Pembahasan.....	77
4.2.1	Pembahasan Hasil Uji Coba Skenario 1.....	77
4.2.2	Pembahasan Hasil Uji Coba Skenario 2.....	79
4.2.3	Pembahasan Hasil Uji Coba Skenario 3.....	80
4.2.4	Pembahasan Hasil Uji Coba Skenario 4.....	82
4.2.5	Pembahasan Antar Skenario	84
4.3	Integrasi Islam	87

4.3.1	<i>Muamalah Ma'a Allah Subhanahu wa Ta'ala</i>	88
4.3.2	<i>Muamalah Ma'a An-Nas</i>	89
BAB V	KESIMPULAN DAN SARAN	92
5.1	Kesimpulan	92
5.2	Saran.....	94
DAFTAR PUSTAKA		
LAMPIRAN		

DAFTAR TABEL

Tabel 2.1 Penelitian Terdahulu	29
Tabel 3.1 Rincian Data dari Platform	36
Tabel 3.2 Proses <i>Case Folding</i>	41
Tabel 3.3 Proses <i>Cleansing</i>	40
Tabel 3.4 Proses <i>Normalization</i>	42
Tabel 3.5 Proses <i>Tokenization</i>	43
Tabel 3.6 Proses <i>Stopword Removal</i>	43
Tabel 3.7 Proses <i>Stemming</i>	45
Tabel 3.8 Contoh Dokumen Hasil <i>Preprocessed</i>	47
Tabel 3.9 Daftar Kosakata (Fitur)	47
Tabel 3.10 Hasil Perhitungan Frekuensi Kata (TF)	48
Tabel 3.11 Hasil Perhitungan <i>Inverse Document Frequency</i> (IDF)	49
Tabel 3.12 Hasil Perhitungan Bobot TF-IDF	50
Tabel 3.13 Cuplikan <i>Dataset</i> Siap Pakai untuk <i>Random Forest</i>	50
Tabel 3.14 <i>Confusion Matrix Aspect</i>	57
Tabel 3.15 <i>Confusion Matrix Sentiment</i>	58
Tabel 3.16 Rasio Pembagian Data	60
Tabel 3.17 Tahapan Klasifikasi Dua Tingkat	61
Tabel 4.1 Confusion Matrix Klasifikasi Aspek Skenario 1	65
Tabel 4.2 Evaluasi Metrik Klasifikasi Aspek Skenario 1	66
Tabel 4.3 Confusion Matrix Klasifikasi Sentimen Skenario 1	67
Tabel 4.4 Evaluasi Metrik Klasifikasi Sentimen Skenario 1	67
Tabel 4.5 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 2	69
Tabel 4.6 Evaluasi Metrik Klasifikasi Aspek Skenario 2	69
Tabel 4.7 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 2	70
Tabel 4.8 Evaluasi Metrik Klasifikasi Sentimen Skenario 2	71
Tabel 4.9 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 3	72
Tabel 4.10 Evaluasi Metrik Klasifikasi Aspek Skenario 3	72
Tabel 4.11 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 3	73
Tabel 4.12 Evaluasi Metrik Klasifikasi Sentimen Skenario 3	73
Tabel 4.13 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 4	75
Tabel 4.14 Evaluasi Metrik Klasifikasi Aspek Skenario 4	75
Tabel 4.15 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 4	76
Tabel 4.16 Evaluasi Metrik Klasifikasi Sentimen Skenario 4	76
Tabel 4.17 Rekapitulasi Performa Klasifikasi Aspek 4 Skenario	85
Tabel 4.18 Rekapitulasi Performa Klasifikasi Sentimen 4 Skenario	85

DAFTAR GAMBAR

Gambar 3.1 Perancangan Alur Penelitian	33
Gambar 3.2 Proses Pengumpulan Data.....	34
Gambar 3.3 Proses Anotasi Data	37
Gambar 3.4 Alur <i>PreProcessing</i> Data.....	39
Gambar 3.5 TF-IDF	46
Gambar 3.6 Simulasi <i>Bootstrap Sampling</i>	52
Gambar 3.7 Ilustrasi Pembentukan Pohon Keputusan	55
Gambar 4.1 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 1.....	65
Gambar 4.2 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 1.....	66
Gambar 4.3 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 2.....	69
Gambar 4.4 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 2.....	70
Gambar 4.5 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 3.....	72
Gambar 4.6 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 3.....	73
Gambar 4.7 <i>Confusion Matrix</i> Klasifikasi Aspek Skenario 4.....	75
Gambar 4.8 <i>Confusion Matrix</i> Klasifikasi Sentimen Skenario 4.....	76

ABSTRAK

Ashidqy, Moch. Hasbi. 2026. Analisis Sentimen Berbasis Aspek pada Ulasan Wisata Jatim Park Group Menggunakan *Random Forest*. Skripsi. Program Studi Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Supriyono, M.Kom (II) Dr. H. Mochamad Imamudin, Lc., M.A

Kata Kunci: Analisis sentimen berbasis aspek, Jatim Park Group, Konsep 3A, Random Forest, TF-IDF

Pariwisata merupakan sektor strategis yang menghasilkan jumlah ulasan digital yang besar, khususnya pada destinasi wisata tematik berskala besar seperti Jatim Park Group di Kota Batu, Jawa Timur. Ulasan tersebut memuat opini wisatawan terhadap berbagai aspek pengalaman berkunjung yang bersifat tidak terstruktur dan beragam, sehingga analisis manual menjadi tidak efisien. Penelitian ini mengusulkan pendekatan *Aspect-Based Sentiment Analysis* (ABSA) menggunakan algoritma *Random Forest* untuk mengklasifikasikan sentimen ulasan wisata Jatim Park Group berdasarkan kerangka konsep 3A, yaitu Atraksi, Aksesibilitas, dan Amenitas, serta menganalisis pengaruh jumlah pohon keputusan ($n_estimators$) terhadap kestabilan performa model. Data yang digunakan sebanyak 4.195 ulasan, diperoleh dari platform Google Maps dan TripAdvisor periode 2020–2025. Ekstraksi fitur teks dilakukan menggunakan metode TF-IDF, sedangkan anotasi data dilakukan secara otomatis menggunakan *Large Language Model* dengan tingkat kesesuaian validasi pakar sebesar 82%. Pengujian dilakukan melalui dua tahap klasifikasi, yaitu Klasifikasi Aspek dan Klasifikasi Sentimen, pada 4 skenario yang memvariasikan rasio pembagian data dan nilai $n_estimators$. Hasil penelitian menunjukkan bahwa konfigurasi terbaik diperoleh pada skenario 4 dengan rasio pembagian data 80:20 dan $n_estimators$ 200, menghasilkan rata-rata F1-score Klasifikasi Aspek sebesar 0,61 dan Klasifikasi Sentimen sebesar 0,94. Performa model mencerminkan karakteristik data ulasan wisata secara nyata bahwa aspek Atraksi mendominasi ulasan pengunjung, tercermin dari tingginya performa model pada aspek tersebut dibandingkan aspek Aksesibilitas dan Amenitas.

ABSTRACT

Ashidqy, Moch. Hasbi. 2026. Analisis Sentimen Berbasis Aspek pada Ulasan Wisata Jatim Park Group Menggunakan *Random Forest*. Skripsi. Program Studi Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Supriyono, M.Kom (II) Dr. H. Mochamad Imamudin, Lc., M.A

Keywords: Aspect-based sentiment analysis, Jatim Park Group, 3A concept, Random Forest, TF-IDF

Tourism represents a strategic sector that generates a substantial volume of digital reviews, particularly for large-scale thematic destinations such as the Jatim Park Group in Batu City, East Java. These reviews contain tourist opinions regarding various aspects of their visiting experiences, which are unstructured and diverse, thereby rendering manual analysis inefficient. This study proposes an Aspect-Based Sentiment Analysis (ABSA) approach utilizing the Random Forest algorithm to classify the sentiment of Jatim Park Group tourism reviews based on the 3A conceptual framework (Attraction, Accessibility, and Amenity), while also analyzing the effect of the number of decision trees ($n_{estimators}$) on the stability of model performance. The dataset comprises 4,195 reviews sourced from Google Maps and TripAdvisor platforms spanning the 2020–2025 period. Text feature extraction was performed using the TF-IDF method, whereas data annotation was executed automatically employing a Large Language Model, achieving an expert validation agreement rate of 82%. Testing was conducted through two independent classification stages: Aspect Classification and Sentiment Classification, across four scenarios that varied the data split ratios and $n_{estimators}$ values. The results demonstrate that the optimal configuration was achieved in scenario 4, with an 80:20 data split ratio and an $n_{estimators}$ value of 200, yielding an average F1-score of 0.61 for Aspect Classification and 0.94 for Sentiment Classification. The model performance aligns with the empirical characteristics of the review data, demonstrating that the dominance of the Attraction aspect within visitor reviews directly corresponds to the higher classification accuracy in that specific aspect compared to the Accessibility and Amenity aspects.

المُلخَص

محمد حسيبي الصديقي. ٢٠٢٦. تحليل المشاعر القائم على الجوانب في مراجعات السياحة لمجموعة جاتيم بارك باستخدام الغابة العشوائية. بحث جامعي. قسم الهندسة المعلوماتية، كلية العلوم والتكنولوجيا، جامعة مولانا مالك إبراهيم الإسلامية الحكومية مالانج. المشرفان: (أولاً) سوبريونو، ماجستير (ثانياً) د. الحاج محمد إمام الدين، ليسانس، ماجستير

الكلمات الرئيسية: تحليل المشاعر المستند إلى الجوانب، مجموعة جاتيم بارك، مفهوم A_3 ، الغابة العشوائية، $TF-IDF$

تُعَدُّ السياحة قطاعاً استراتيجياً يُنتج حجماً كبيراً من المراجعات الرقمية، لا سيما في الوجهات السياحية الموضوعية الكبرى كمجموعة جاتيم بارك في مدينة باتو، جاوى الشرقية. تحتوي هذه المراجعات على آراء السياح حول جوانب متعددة من تجربة الزيارة، وهي ذات طابع غير منظم ومتنوع، مما يجعل التحليل اليدوي غير فعال. تقترح هذه الدراسة نهج تحليل المشاعر القائم على الجوانب (ABSA) باستخدام خوارزمية الغابة العشوائية (*Random Forest*) لتصنيف مشاعر مراجعات سياحة مجموعة جاتيم بارك وفق إطار مفهوم A_3 ، وهي الجاذبية والإمكانية والمرافق، فضلاً عن تحليل تأثير عدد أشجار القرار ($n_estimators$) على استقرار أداء النموذج. تتكون البيانات المستخدمة من ٤,١٩٥ مراجعة مستخرجة من منصتي *Google Maps* و *TripAdvisor* للفترة الممتدة من ٢٠٢٠ إلى ٢٠٢٥. جرى تمثيل ميزات النص باستخدام طريقة $TF-IDF$ ، فيما نُقِدَ التوصيف التلقائي للبيانات بواسطة نموذج لغوي كبير (*Large Language Model*) بمعدل توافق مع تحقق الخبراء بلغ ٨٢٪. أُجري الاختبار عبر مرحلتين مستقلتين من التصنيف، هما تصنيف الجوانب وتصنيف المشاعر، وذلك عبر أربعة سيناريوهات تتباين في نسبة تقسيم البيانات وقيمة $n_estimators$. أظهرت النتائج أن أفضل تهئية تحققت في السيناريو الرابع بنسبة تقسيم بيانات ٨٠:٢٠ وقيمة $n_estimators$ تساوي ٢٠٠، محققةً متوسط $F1-score$ بلغ ٠,٦١ لتصنيف الجوانب و ٠,٩٤ لتصنيف المشاعر. يعكس أداء النموذج الخصائص الطبيعية لبيانات مراجعات السياحة، إذ تجلّت هيمنة جانب الجاذبية في مراجعات السياح من خلال ارتفاع أداء النموذج على هذا الجانب مقارنةً بجانب إمكانية الوصول والمرافق

BAB I

PENDAHULUAN

1.1 Latar belakang

Pariwisata menjadi sektor strategis yang menyumbang dampak besar terhadap kemajuan ekonomi di skala lokal dan nasional. Aktivitas pariwisata tidak hanya berdampak pada peningkatan pendapatan dan penciptaan lapangan kerja, tetapi juga membentuk persepsi serta pengalaman wisatawan terhadap suatu destinasi (Maulidiya Indah Sari, Masrurotun Mufidah, Sarpini, 2023). Dalam era digital, pengalaman tersebut tidak lagi berhenti pada kunjungan fisik saja melainkan terdokumentasi dalam bentuk ulasan daring di berbagai platform digital. Ulasan ini menjadi representasi opini publik yang dapat dimanfaatkan sebagai sumber informasi dan evaluasi kualitas destinasi wisata (Ariansyah et al., 2025).

Terletak dalam cakupan wilayah Malang Raya, Kota Batu dikenal sebagai pusat pariwisata unggulan di Jawa Timur. Keunggulan kompetitif daerah ini terletak pada objek wisata tematik yang mengintegrasikan nilai edukatif serta rekreatif bagi pengunjung (Kecamatan Batu, n.d.). Data yang dirilis oleh otoritas pariwisata Kota Batu mempresentasikan besarnya arus wisatawan pada 2025 yang terakumulasi sebanyak kurang lebih 9,7 juta orang (Andre Setiawan, 2025). Selain itu, pada periode libur Lebaran 2025 tercatat lebih dari 778 ribu wisatawan mengunjungi berbagai objek wisata di Kota Batu dalam kurun waktu libur yang relatif singkat (Pradana, 2025). Tingginya angka kunjungan ini menunjukkan jumlah interaksi wisatawan yang besar terhadap destinasi wisata di Kota Batu dan

secara tidak langsung berimplikasi pada meningkatnya jumlah ulasan digital yang dihasilkan.

Dominasi sektor pariwisata di Kota Batu tidak lepas dari peran Jatim Park Group dalam menginisiasi destinasi wisata tematik. Fokusnya pada integrasi edukasi dan hiburan keluarga terwujud melalui sejumlah objek wisata kenamaan, di antaranya adalah Jatim Park seri 1 hingga 3, Museum Angkut, dan unit rekreasi Batu Night Spectacular. Keberagaman wahana serta segmentasi pasar keluarga menjadikan Jatim Park Group sebagai tujuan wisata modern di Kota Batu dengan tingkat kunjungan yang tinggi sepanjang tahun, khususnya pada musim liburan dan hari besar nasional (Jatim Park Group, 2024). Skala operasional yang besar dan kompleksitas layanan yang beragam menyebabkan destinasi ini menghasilkan volume ulasan daring yang signifikan pada berbagai platform seperti *Google Maps* dan situs perjalanan lainnya.

Ulasan-ulasan tersebut memuat opini wisatawan terhadap berbagai aspek pengalaman berkunjung mulai dari kualitas wahana, fasilitas pendukung, kebersihan, pelayanan, hingga kemudahan akses dan manajemen antrean. Namun demikian, ulasan wisata umumnya berbentuk teks tidak terstruktur dengan gaya bahasa yang beragam, termasuk bahasa informal dan opini subjektif (Asyiah & Aktavia, 2025). Satu ulasan bahkan dapat memuat lebih dari satu sentimen terhadap aspek yang berbeda (Arifiyanti et al., 2022). Kondisi ini menyebabkan analisis manual menjadi tidak efisien serta berpotensi mengabaikan informasi penting yang tersembunyi di dalam data teks.

Penelitian mengenai analisis sentimen berbasis aspek telah banyak dikembangkan, salah satunya oleh Firdaus dkk. (2025) yang mengkaji ulasan aplikasi Alfagift menggunakan algoritma *Random Forest*. Studi tersebut membuktikan bahwa *Random Forest* sangat efektif dalam menangani data teks yang kompleks dengan tingkat akurasi mencapai 89%. Keberhasilan penggunaan *Random Forest* dalam mengklasifikasikan sentimen secara akurat menjadi landasan utama bagi penelitian ini untuk menerapkan algoritma serupa. Namun, berbeda dengan penelitian tersebut yang berfokus pada domain aplikasi belanja daring, penelitian ini akan menguji ketangguhan metode tersebut pada domain pariwisata yang memiliki karakteristik ulasan yang lebih bervariasi.

Dalam kajian pariwisata, kualitas destinasi umumnya dianalisis menggunakan konsep 3A, yaitu *Attraction* (daya tarik wisata), *Amenity* (fasilitas penunjang), dan *Accessibility* (kemudahan akses). Ketiga elemen ini berfungsi sebagai pilar evaluasi dalam menentukan sejauh mana sebuah destinasi mampu menciptakan kepuasan bagi para wisatawan (Seran et al., 2023). Atas dasar tersebut, penelitian ini mengadopsi teknik analisis sentimen yang mampu membedah opini pengunjung ke dalam dimensi aspek 3A. Pendekatan ini melampaui klasifikasi sentimen sederhana karena memberikan rincian data yang lebih akurat mengenai kondisi destinasi.

Karakteristik ulasan yang beragam dan tidak terstruktur menuntut sebuah metodologi yang mampu mengurai opini wisatawan secara objektif agar tidak terjadi bias dalam interpretasi. Keharusan untuk melakukan penyaringan informasi secara teliti ini bukan sekadar kebutuhan teknis dalam analisis data, melainkan juga

merupakan bentuk integritas dalam menyampaikan fakta. Perspektif Islam memandang bahwa memverifikasi setiap detail ulasan adalah sebuah keharusan demi menjaga integritas dan kejujuran informasi, sebagaimana prinsip ketelitian dalam menerima berita. Oleh karena itu, urgensi dalam mengolah ulasan yang bervariasi dan tidak terstruktur ini pada hakikatnya merupakan implementasi dari prinsip *Tabayyun* sebagaimana diamanatkan dalam Al-Qur'an Surat Al-Hujurat (49:6) :

يَا أَيُّهَا الَّذِينَ آمَنُوا إِنْ جَاءَكُمْ فَاسِقٌ بِنَبَأٍ فَتَبَيَّنُوا أَنْ تُصِيبُوا قَوْمًا بِجَهَالَةٍ فَتُصْبِحُوا عَلَىٰ مَا فَعَلْتُمْ نَادِمِينَ

“Wahai orang-orang yang beriman, jika seorang fasik datang kepadamu membawa berita penting, maka telitilah kebenarannya agar kamu tidak mencelakakan suatu kaum karena ketidaktahuan(-mu) yang berakibat kamu menyesali perbuatanmu itu.” (QS. Al-Hujurat [49]: 6)

Ayat ini menegaskan pentingnya ketelitian dalam memverifikasi sebuah informasi guna menghindari kesimpulan yang keliru (Lajnah Pentashihan Mushaf Al-Qur'an, 2019). Dalam Tafsir Ibnu Katsir, Ibnu Katsir menjelaskan bahwa perintah *tabayyun* dalam ayat ini bertujuan agar kaum mukmin tidak menimpakan musibah kepada suatu kaum tanpa mengetahui keadaan yang sebenarnya, sejalan dengan sabda Nabi Muhammad yang dikutip dalam tafsir tersebut bahwa kehati-hatian itu datang dari Allah sedangkan ketergesa-gesaan datang dari setan (Al-Sheikh, 2004).

Dalam konteks penelitian ini, ulasan wisatawan yang bersifat subjektif dan beragam diperlakukan sebagai 'berita' yang memerlukan proses verifikasi mendalam melalui *Aspect-Based Sentiment Analysis* (ABSA). Penggunaan algoritma *Random Forest* berfungsi sebagai instrumen teknis untuk melakukan

klasifikasi yang objektif, yang secara praktis memanifestasikan perintah *tabayyun* tersebut. Untuk mengoperasionalkan prinsip ketelitian ini dalam ranah pariwisata, penelitian ini mengadopsi konsep 3A, yaitu *Attraction* (daya tarik), *Amenity* (fasilitas), dan *Accessibility* (aksesibilitas). Ketiga elemen ini berfungsi sebagai pilar evaluasi dalam menentukan kepuasan wisatawan (Seran et al., 2023). Dengan membedah ulasan ke dalam bagian-bagian terkecil, penelitian ini memastikan bahwa setiap opini ditempatkan sesuai proporsinya secara adil, sehingga manajemen destinasi tidak melakukan generalisasi yang dapat merugikan salah satu aspek pelayanan. Pendekatan sistematis ini menjadi jembatan untuk mengevaluasi kualitas destinasi secara lebih komprehensif.

Berdasarkan uraian tersebut, tingginya angka kunjungan wisata di Kota Batu serta masifnya volume ulasan digital terhadap destinasi Jatim Park Group menyajikan potensi data besar yang memerlukan metode analisis secara sistematis dan objektif. Upaya untuk menjaga ketelitian informasi serta menerapkan prinsip *Tabayyun* dalam mengolah data teks mendasari penelitian ini untuk mengimplementasikan pendekatan Analisis Sentimen Berbasis Aspek (ABSA) menggunakan algoritma *Random Forest* berdasarkan kerangka 3A. Temuan yang dihasilkan diharapkan dapat menjadi referensi dalam pengembangan metode klasifikasi teks di sektor pariwisata serta membuktikan efektivitas analisis sentimen sebagai instrumen evaluasi yang akurat berbasis data.

1.2 Rumusan Masalah

1. Bagaimana performa metode *Random Forest* dalam mengklasifikasikan sentimen berbasis aspek pada ulasan wisata Jatim Park Group?

2. Bagaimana pengaruh jumlah pohon keputusan ($n_estimators$) pada *Random Forest* terhadap performa klasifikasi ulasan?

1.3 Tujuan Penelitian

1. Mengklasifikasikan sentimen ulasan Jatim Park Group berbasis aspek 3A menggunakan metode *Random Forest* serta mengukur tingkat akurasi, presisi, *recall*, dan *F1-score*-nya.
2. Menganalisis pengaruh perubahan jumlah pohon ($n_estimators$) pada algoritma *Random Forest* untuk mendapatkan model dengan performa paling stabil.

1.4 Batasan Penelitian

- a. Data yang digunakan berupa ulasan wisata Jatim Park Group yang diperoleh dari platform *Google Maps* dan *TripAdvisor* pada periode 2020 -2025.
- b. Analisis sentimen dilakukan menggunakan pendekatan *Aspect-Based Sentiment Analysis* dengan pembagian aspek berdasarkan konsep 3A, yaitu *Attraction*, *Amenity*, dan *Accessibility*.
- c. Kajian ini secara spesifik menerapkan algoritma *Random Forest* sebagai metode klasifikasi tunggal tanpa melibatkan komparasi dengan model pembelajaran mesin lainnya.
- d. Penilaian terhadap kinerja model diukur melalui metrik evaluasi yang meliputi *F1-score*, *recall*, *precision*, serta *accuracy*.

- e. Ruang lingkup studi ini dibatasi pada pengujian performa model klasifikasi, sehingga tidak mencakup pembahasan mengenai strategi manajerial dalam tata kelola objek wisata secara ekstensif.

1.5 Manfaat Penelitian

1. Manfaat Teoritis

Penelitian ini diharapkan dapat memberikan kontribusi dalam pengembangan kajian *text mining* dan analisis sentimen, khususnya pada penerapan *Aspect-Based Sentiment Analysis* (ABSA) dalam domain pariwisata. Selain itu, penelitian ini dapat menjadi referensi empiris mengenai performa metode *Random Forest* dalam mengklasifikasikan sentimen berbasis aspek pada data teks tidak terstruktur.

2. Manfaat Praktisi

Secara praktis, penelitian ini dapat memberikan gambaran mengenai kemampuan metode *Random Forest* dalam mengolah dan mengklasifikasikan ulasan wisata berbasis aspek 3A (*Attraction, Amenity, Accessibility*). Bagi pengelola destinasi wisata, khususnya Jatim Park Group, hasil penelitian ini dapat menjadi dasar dalam mempertimbangkan penggunaan sistem analisis sentimen otomatis untuk memantau opini wisatawan secara lebih efisien.

BAB II

STUDI PUSTAKA

2.1 Konsep 3A dalam Pariwisata

Pembangunan dan pengembangan destinasi pariwisata yang kompetitif memerlukan keterpaduan antara berbagai elemen pendukung yang membentuk satu kesatuan produk wisata. Secara teoretis, komponen-komponen utama yang membentuk suatu destinasi diringkas dalam konsep 3A yang terdiri dari Atraksi, Aksesibilitas dan Amenitas. Menurut Stange, Jennifer dan Brown, David (2011) dalam *Tourism Destination Management: Achieving Sustainable and Competitive Results*, keberhasilan suatu destinasi sangat ditentukan oleh kemampuan pengelola dalam mengintegrasikan ketiga aspek tersebut secara sistematis dan berkelanjutan. Ketiganya tidak berdiri sendiri, melainkan saling mendukung dalam membentuk pengalaman wisata yang utuh (Stange et al., 2010). Dalam penelitian ini, ketiga komponen tersebut dijadikan dasar penentuan kategori aspek dalam proses anotasi data ulasan wisatawan, sehingga analisis sentimen yang dihasilkan dapat mencerminkan dimensi pengalaman wisata secara terstruktur.

Atraksi merupakan komponen utama yang mendorong seseorang untuk memutuskan berkunjung ke suatu destinasi wisata dan sering kali menjadi pertimbangan pertama wisatawan sebelum menentukan pilihan perjalanan (Safari et al., 2025). Suatu wilayah akan sulit berkembang apabila tidak memiliki daya tarik yang jelas dan berbeda dari tempat lain, sehingga atraksi berperan sebagai magnet utama yang membentuk citra dan identitas destinasi. Secara umum, daya tarik dapat

dikelompokkan ke dalam tiga kategori utama yaitu daya tarik alam yang mencakup pantai, pegunungan, dan bentang alam unik; daya tarik budaya yang bersumber dari tradisi, kesenian, adat istiadat, hingga warisan sejarah; serta daya tarik buatan atau *man-made attraction* seperti taman rekreasi, museum interaktif dan kawasan wisata tematik (Supatmana & Suwarti, 2022). Keberagaman daya tarik inilah yang pada dasarnya menggerakkan manusia untuk berwisata, sebagaimana diisyaratkan dalam QS. Al-Ankabut ayat 20:

قُلْ سِيرُوا فِي الْأَرْضِ فَانظُرُوا كَيْفَ بَدَأَ الْخَلْقَ ثُمَّ اللَّهُ يُنشِئُ النَّشْأَةَ الْآخِرَةَ إِنَّ اللَّهَ عَلَىٰ كُلِّ شَيْءٍ قَدِيرٌ

"Katakanlah, "Berjalanlah di (muka) bumi, lalu perhatikanlah bagaimana Allah memulai penciptaan (semua makhluk). Kemudian, Allah membuat kejadian yang akhir (setelah mati di akhirat kelak). Sesungguhnya Allah Mahakuasa atas segala sesuatu." (QS. Al-Ankabut[29]: 20)

Dari ayat tersebut mengandung makna bahwa perjalanan untuk menyaksikan keunikan suatu tempat merupakan bagian dari fitrah manusia. Dalam konteks Jatim Park Group sebagai kawasan wisata tematik buatan, aspek atraksi menjadi dimensi yang paling dominan dibicarakan dalam ulasan wisatawan sehingga menjadi salah satu kategori aspek utama dalam penelitian ini.

Aksesibilitas dalam konteks pariwisata merujuk pada tingkat kemudahan yang dimiliki wisatawan untuk mencapai suatu destinasi serta bergerak secara nyaman di dalam kawasan tersebut (Safari et al., 2025). Konsep ini tidak hanya berkaitan dengan jarak tempuh, tetapi juga menyangkut kualitas sarana dan prasarana yang tersedia. Mulai dari infrastruktur transportasi sebagai fondasi utama hingga sistem informasi yang membantu wisatawan merencanakan perjalanan. Secara lebih rinci, aksesibilitas dapat dibedakan menjadi beberapa dimensi yang

saling melengkapi yaitu akses eksternal yang mencakup ketersediaan jalan raya, bandara, dan transportasi umum; akses internal yang berkaitan dengan kemudahan mobilitas di dalam area destinasi seperti rambu penunjuk arah dan jalur pedestrian; serta akses informasi yang semakin relevan di era digital melalui ketersediaan peta, situs resmi, dan media promosi daring (Apriani et al., 2025). Prinsip kemudahan dalam aksesibilitas ini sejalan dengan sabda Nabi Muhammad *Shallallahu 'Alaihi wa Sallam*:

يَسِّرُوا وَلَا تُعَسِّرُوا، وَبَشِّرُوا وَلَا تُنْفِرُوا

"Mudahkanlah dan jangan mempersulit, berilah kabar gembira dan jangan membuat orang lari." (HR. Bukhari No. 69)

Dari Hadis tersebut mengisyaratkan bahwa menciptakan kemudahan bagi orang lain merupakan nilai yang dianjurkan. Hambatan pada salah satu dimensi tersebut dapat menimbulkan persepsi negatif meskipun destinasi memiliki potensi besar, sehingga ulasan wisatawan terkait aksesibilitas menjadi sinyal penting yang perlu diidentifikasi dalam analisis sentimen berbasis aspek pada penelitian ini.

Amenitas merupakan komponen pendukung yang berperan dalam memastikan wisatawan merasa nyaman selama berada di suatu destinasi (Seran et al., 2023). Aspek ini bukan menjadi alasan utama seseorang memilih tempat berkunjung, kualitas amenities sangat memengaruhi kesan yang ditinggalkan karena kekurangannya dapat menurunkan tingkat kepuasan meskipun daya tarik destinasi cukup kuat. Komponen amenities mencakup berbagai kategori fasilitas yang saling melengkapi dari akomodasi seperti hotel dan penginapan untuk memenuhi kebutuhan istirahat, fasilitas konsumsi berupa restoran, kafe, dan pusat kuliner,

hingga fasilitas umum seperti toilet, area parkir, pusat informasi dan tempat ibadah (Leylita Novita Rossadi & Endang Widayati, 2024). Selain itu, layanan pendukung seperti pemandu wisata, pusat oleh-oleh, serta layanan kesehatan darurat turut berkontribusi dalam memberikan rasa aman dan kemudahan selama beraktivitas (Seran et al., 2023). Semangat pelayanan ini sejalan dengan sabda Nabi Muhammad *Shallallahu 'Alaihi wa Sallam*:

مَنْ كَانَ يُؤْمِنُ بِاللَّهِ وَالْيَوْمِ الْآخِرِ فَلْيُكْرِمْ ضَيْفَهُ

"Barangsiapa beriman kepada Allah dan hari akhir, maka muliakanlah tamunya."
(HR. Bukhari No. 6018)

Hadis tersebut menegaskan bahwa memuliakan tamu termasuk memastikan kenyamanan dan kelengkapan fasilitas yang mereka butuhkan merupakan cerminan nilai keimanan. Kelengkapan dan kualitas fasilitas tersebut sering kali tercermin secara eksplisit dalam ulasan wisatawan, sehingga amenitas menjadi kategori aspek yang relevan untuk diidentifikasi dalam analisis sentimen pada penelitian ini.

Ketiga komponen dalam konsep 3A yaitu atraksi, aksesibilitas, dan amenitas secara bersama-sama membentuk dimensi pengalaman wisatawan yang dapat dievaluasi secara terstruktur. Dalam konteks analisis sentimen berbasis aspek, konsep ini memberikan landasan yang sistematis untuk mengelompokkan opini wisatawan ke dalam kategori yang bermakna dan dapat diinterpretasikan secara kontekstual. Setiap komponen merepresentasikan dimensi pengalaman yang berbeda namun saling berkaitan, sehingga sentimen yang teridentifikasi pada masing-masing aspek dapat memberikan gambaran evaluatif yang lebih spesifik dan informatif dibandingkan analisis sentimen secara keseluruhan. Oleh karena itu,

dalam penelitian ini konsep 3A diadopsi sebagai dasar penentuan kategori aspek pada proses anotasi data ulasan wisatawan Jatim Park Group, dengan tujuan menghasilkan model klasifikasi sentimen berbasis aspek yang mencerminkan dimensi kualitas destinasi secara menyeluruh.

2.2 Analisis Sentimen Berbasis Aspek

Aspect-Based Sentiment Analysis (ABSA) merupakan pendekatan lanjutan dalam analisis sentimen yang bertujuan untuk mengidentifikasi aspek spesifik dari suatu entitas serta menentukan polaritas sentimen terhadap masing-masing aspek tersebut (Pontiki et al., 2014). Berbeda dengan analisis sentimen pada tingkat dokumen yang hanya menghasilkan satu label sentimen secara keseluruhan, ABSA memungkinkan identifikasi opini secara lebih terperinci dengan memetakan sentimen terhadap komponen tertentu yang dibahas dalam teks. Pendekatan ini banyak digunakan dalam analisis ulasan daring karena satu ulasan sering kali mengandung beberapa aspek dengan polaritas yang berbeda.

Secara konseptual, *Aspect-Based Sentiment Analysis* (ABSA) berfokus pada dua komponen utama, yaitu aspek dan polaritas sentimen. Aspek merujuk pada atribut, komponen atau karakteristik tertentu dari suatu entitas yang menjadi objek evaluasi dalam teks. Polaritas sentimen menunjukkan sikap atau opini penulis terhadap aspek tersebut yang umumnya dikategorikan sebagai positif, negatif atau netral. Dengan memisahkan kedua komponen ini, ABSA mampu memberikan analisis yang lebih terperinci dibandingkan analisis sentimen konvensional pada tingkat dokumen.

Sebagai contoh, dalam kalimat "Wahana bermainnya sangat seru dan beragam untuk semua usia, fasilitas toilet dan musholanya kurang bersih dan antreannya sangat panjang, tetapi lokasi parkir kendaraannya luas dan mudah dijangkau dari pusat kota" terdapat tiga aspek berbeda sesuai konsep 3A pariwisata, yaitu atraksi, amenitas, dan aksesibilitas. Masing-masing aspek memiliki polaritas sentimen yang tidak sama: atraksi berupa wahana hiburan dinilai positif, amenitas berupa fasilitas umum dinilai negatif, sedangkan aksesibilitas berupa kemudahan parkir dan lokasi dinilai positif. Pendekatan ABSA (*Aspect-Based Sentiment Analysis*) memungkinkan identifikasi setiap aspek beserta sentimennya secara terpisah dalam satu kalimat yang sama.

Secara umum, proses dalam ABSA mencakup tahap identifikasi aspek yang relevan dalam teks dan penentuan sentimen terhadap masing-masing aspek tersebut. Dalam praktiknya, aspek dapat muncul secara eksplisit ketika disebutkan langsung dalam kalimat, maupun secara implisit ketika disimpulkan dari konteks. Kemampuan untuk menangani aspek implisit memerlukan pemahaman konteks yang lebih mendalam sehingga meningkatkan kompleksitas analisis.

Dalam kajian *Aspect-Based Sentiment Analysis* (ABSA), banyak penelitian merujuk pada kerangka evaluasi yang diperkenalkan melalui *SemEval 2014 Task 4* oleh Pontiki et al. sebagai dasar metodologis. Kerangka tersebut membagi proses analisis ke dalam beberapa subtugas yang terstruktur dan sistematis. Meskipun demikian, tidak seluruh subtugas harus diterapkan secara bersamaan dalam setiap penelitian. Peneliti dapat memilih tahapan yang paling relevan dengan tujuan dan ruang lingkup kajian yang dilakukan. Dalam konteks penelitian ini, fokus diarahkan

pada tahap kategorisasi aspek serta klasifikasi polaritas sentimen. Pendekatan tersebut tetap berada dalam koridor metodologi ABSA yang telah teruji dalam berbagai studi sebelumnya. Penyesuaian ini dilakukan agar analisis lebih terarah dan sesuai dengan karakteristik data yang digunakan. Dengan pembatasan tersebut, pembahasan teoritis dapat difokuskan pada aspek yang mendukung proses klasifikasi secara mendalam.

Subtask 3 dalam kerangka tersebut dikenal sebagai *Aspect Category Detection* (ACD). Secara konseptual, ACD bertujuan mengidentifikasi kategori aspek yang terkandung dalam suatu teks berdasarkan skema yang telah ditentukan sebelumnya. Berbeda dengan pendekatan ekstraksi istilah aspek yang berorientasi pada pencarian frasa spesifik, ACD menitikberatkan pada pengelompokan opini ke dalam kategori konseptual. Dalam penelitian ini, kategori aspek disusun mencakup atraksi, aksesibilitas dan amenitas. Pengelompokan ini bertujuan untuk menyederhanakan keragaman ekspresi dalam ulasan wisata menjadi struktur yang lebih sistematis. Model klasifikasi kemudian dikembangkan untuk mengenali pola bahasa yang mengindikasikan keterkaitan dengan masing-masing kategori. Ketepatan pada tahap deteksi kategori menjadi fondasi penting bagi analisis lanjutan. Hasil identifikasi kategori akan menentukan konteks evaluasi sentimen terhadap setiap aspek destinasi.

Subtask 4 dikenal sebagai *Aspect Sentiment Polarity Classification* (APC). Tahap ini berfokus pada penentuan kecenderungan sentimen terhadap kategori aspek yang telah diidentifikasi sebelumnya. Secara teoritis, proses ini mengklasifikasikan opini ke dalam polaritas seperti positif, negatif, atau netral

berdasarkan konteks pernyataan. Implementasinya umumnya melibatkan teknik representasi fitur teks serta algoritma pembelajaran mesin. Dalam penelitian ini, metode *Random Forest* digunakan sebagai pendekatan klasifikasi untuk memodelkan hubungan antara fitur teks dan label sentimen. Kinerja model selanjutnya dievaluasi menggunakan metrik seperti *accuracy*, *precision*, *recall*, dan *F1-score*. Pembatasan pada *Subtask 3* dan *Subtask 4* memungkinkan penelitian tetap konsisten dengan kerangka ABSA yang baku sehingga analisis opini wisatawan dapat dilakukan secara terstruktur tanpa melalui tahap ekstraksi istilah aspek secara eksplisit.

2.3 *Preprocessing Data*

Preprocessing data merupakan tahapan awal yang krusial dalam *pipeline Natural Language Processing* (NLP), khususnya pada tugas *Aspect-Based Sentiment Analysis* (ABSA). Tahap ini bertujuan untuk membersihkan, menormalkan, dan mentransformasikan data teks mentah menjadi representasi yang lebih terstruktur sehingga dapat diproses secara optimal oleh model pembelajaran mesin maupun *deep learning* (Savila et al., 2026). Tahapan *preprocessing* pada penelitian ini meliputi:

2.3.1 *Cleansing*

Cleansing merupakan tahap pembersihan teks dari berbagai karakter yang tidak relevan dengan kebutuhan analisis (Aufar et al., 2023). Karakter yang dihapus dapat berupa tanda baca, angka, URL, emoji, maupun simbol khusus lainnya (Supriyono et al., 2024). Keberadaan elemen-elemen tersebut sering kali

menimbulkan noise dalam proses pemodelan. Proses pembersihan bertujuan untuk mempertahankan hanya informasi yang memiliki kontribusi terhadap analisis sentimen. Dengan data yang lebih bersih, kualitas representasi fitur menjadi lebih optimal. Tahap ini menjadi fondasi penting sebelum teks diproses lebih lanjut.

2.3.2 Case Folding

Case folding merupakan proses menyeragamkan seluruh huruf dalam teks menjadi huruf kecil (lowercase). Langkah ini dilakukan untuk menghindari perbedaan token akibat variasi penggunaan huruf kapital (Savila et al., 2026). Dalam pengolahan teks, sistem komputasi dapat membedakan kata yang sama hanya karena perbedaan format penulisan. Oleh sebab itu, normalisasi huruf diperlukan agar representasi data menjadi konsisten. Sebagai contoh, kata “Bagus” dan “bagus” akan diperlakukan sebagai token yang sama setelah melalui proses ini. Dengan demikian, *case folding* membantu mengurangi redundansi kosakata dalam *dataset*.

2.3.3 Normalization

Normalization merupakan tahap awal dalam pemrosesan data teks yang bertujuan untuk menyeragamkan format dokumen agar menjadi lebih konsisten (Sutriawan et al., 2025). Dalam analisis sentimen, data ulasan sering kali memiliki variasi penulisan yang tidak teratur, penggunaan huruf kapital yang acak, serta simbol-simbol yang tidak memiliki nilai semantik (Supriyono et al., 2024). Proses normalisasi memastikan bahwa setiap kata memiliki representasi tunggal yang dapat dikenali secara seragam oleh model klasifikasi.

2.3.4 *Tokenization*

Tokenization adalah proses memecah teks menjadi unit-unit kata yang disebut *token* (Duei Putri et al., 2022). Proses ini memungkinkan teks panjang diuraikan menjadi elemen yang lebih kecil dan terstruktur (Supriyono et al., 2026). Sebagian besar algoritma *Natural Language Processing* bekerja pada tingkat *token*, bukan kalimat utuh. Oleh karena itu, tokenisasi menjadi langkah krusial dalam *pipeline* analisis teks. Hasil dari tahap ini akan digunakan pada proses selanjutnya seperti penghapusan *stopword* dan pembobotan kata. Dengan *tokenisasi* yang tepat, struktur data menjadi lebih mudah dianalisis oleh model.

2.3.5 *Stopword Removal*

Stopword removal bertujuan menghilangkan kata-kata umum yang sering muncul tetapi memiliki makna semantik rendah (Angelina et al., 2023). Contohnya adalah kata seperti “dan”, “yang”, atau “di” dalam bahasa Indonesia. Kata-kata tersebut biasanya tidak memberikan kontribusi signifikan terhadap penentuan sentimen (Supriyono et al., 2026). Jika tidak dihapus, keberadaannya dapat meningkatkan dimensi fitur tanpa menambah informasi penting. Dengan mengurangi kata yang tidak relevan, proses pembelajaran model menjadi lebih efisien. Tahap ini membantu memfokuskan analisis pada kata-kata yang lebih informatif.

2.3.6 *Stemming*

Stemming merupakan proses mengubah kata berimbuhan menjadi bentuk dasarnya. Proses ini membantu menyederhanakan variasi morfologis sehingga kata

dengan akar yang sama tidak dianggap berbeda (Manullang et al., 2023). Sebagai contoh, kata “bermain”, “memainkan”, dan “permainan” dapat direduksi menjadi bentuk dasar yang seragam. Penyederhanaan ini berkontribusi pada pengurangan dimensi fitur dalam *dataset*. Dengan demikian, *stemming* mendukung efisiensi serta konsistensi dalam tahap representasi teks.

Dalam konteks ABSA, *preprocessing* memiliki peran tambahan yaitu mempertahankan informasi aspek yang relevan agar tidak terhapus dalam proses normalisasi. Oleh karena itu, pemilihan teknik *preprocessing* harus disesuaikan dengan karakteristik data dan tujuan analisis.

2.4 *Term Frequency–Inverse Document Frequency (TF-IDF)*

Term Frequency–Inverse Document Frequency (TF-IDF) merupakan metode pembobotan statistik yang umum digunakan dalam *text mining*, *information retrieval*, dan klasifikasi teks untuk mentransformasikan data teks ke dalam representasi numerik (Sutriawan et al., 2024). Pendekatan ini bekerja dengan cara mengukur tingkat kepentingan suatu kata (*term*) dalam sebuah dokumen tertentu secara relatif terhadap keseluruhan dokumen dalam korpus. Secara prinsip, TF-IDF memberikan bobot lebih besar pada kata-kata yang memiliki informasi spesifik dalam sebuah dokumen dan memberikan bobot lebih rendah pada kata umum yang sering muncul di banyak dokumen. Dalam ekosistem *Natural Language Processing* (NLP), teknik ini tetap menjadi standar representasi fitur yang efisien dan efektif, terutama dalam tugas klasifikasi teks sebelum berkembangnya metode berbasis *word embedding* maupun *transformer* (Sutriawan et al., 2025). Secara umum, TF-

IDF terdiri atas dua komponen utama, yaitu Term Frequency (TF) dan Inverse Document Frequency (IDF).

2.4.1 *Term Frequency (TF)*

Term Frequency (TF) digunakan untuk mengukur frekuensi kemunculan suatu kata atau *term* (t) di dalam sebuah dokumen (d) (Rofiqi et al., 2019). Dasar pemikiran dari komponen ini adalah semakin sering suatu kata muncul dalam sebuah dokumen, maka semakin besar pula tingkat relevansi kata tersebut terhadap isi dokumen yang bersangkutan (bobot lokal).

Namun, frekuensi kemunculan mentah (*raw count*) seringkali memiliki kelemahan jika terdapat perbedaan panjang dokumen yang signifikan, terkadang dokumen yang lebih panjang cenderung memiliki frekuensi kata yang lebih tinggi meskipun tingkat relevansinya sama dengan dokumen yang lebih pendek (Septiani & Isabela, 2022). Oleh karena itu, diperlukan proses normalisasi. Secara matematis, persamaan TF dapat dirumuskan sebagai berikut:

$$TF(t, d) = \frac{f(t, d)}{\sum_{t' \in d} f(t', d)} \quad (2.1)$$

keterangan:

- $TF(t, d)$: Nilai *Term Frequency* untuk kata dalam dokumen .
- $f(t, d)$: Jumlah kemunculan term t dalam dokumen d
- $\sum_{t' \in d} f(t', d)$: Total seluruh term dalam dokumen tersebut.

2.4.2 *Inverse Document Frequency (IDF)*

Inverse Document Frequency (IDF) merupakan komponen yang berfungsi untuk mengukur tingkat kelangkaan atau keunikan suatu kata (*term*) di dalam keseluruhan korpus (kumpulan dokumen) (Rofiqi et al., 2019). Berbeda dengan TF

yang berfokus pada frekuensi lokal, IDF berfokus pada signifikansi global sebuah kata. Prinsip utama dari IDF adalah bahwa kata yang muncul di hampir semua dokumen (seperti kata hubung atau kata umum) dianggap memiliki daya pembeda (*discriminative power*) yang rendah, sehingga harus diberikan bobot yang kecil.

Sebaliknya, kata yang hanya muncul pada sedikit dokumen dianggap lebih spesifik dan mampu merepresentasikan topik unik, sehingga diberikan bobot yang tinggi. Secara matematis, penggunaan fungsi logaritma dalam rumus IDF bertujuan untuk menekan efek dari frekuensi dokumen yang sangat besar agar tidak mendominasi perhitungan bobot secara linear (Sutriawan et al., 2025). Secara matematis, persamaan IDF dirumuskan sebagai berikut:

$$IDF(t) = \log \frac{N}{df(t)} \quad (2.2)$$

keterangan:

- N : Jumlah total dokumen dalam korpus,
- $df(t)$: Jumlah dokumen yang mengandung term t .

2.4.3 Perhitungan Bobot TF-IDF

Perhitungan bobot TF-IDF dilakukan untuk mengetahui tingkat kepentingan suatu term dalam sebuah dokumen relatif terhadap keseluruhan dokumen dalam korpus (Rofiqi et al., 2019). Metode ini mengombinasikan nilai *Term Frequency* (TF) yang merepresentasikan frekuensi kemunculan term dalam dokumen dengan *Inverse Document Frequency* (IDF) yang menunjukkan tingkat kelangkaan term pada seluruh dokumen. Semakin sering suatu term muncul dalam dokumen tertentu dan semakin jarang muncul pada dokumen lain, maka bobotnya akan semakin tinggi. Sebaliknya, term yang sering muncul di banyak dokumen akan memiliki

nilai IDF yang rendah sehingga bobot akhirnya juga menurun. Bobot akhir suatu term dalam dokumen diperoleh dengan mengalikan kedua komponen tersebut. Hasil perhitungan ini kemudian digunakan sebagai representasi numerik teks dalam proses analisis lanjutan seperti klasifikasi atau pemodelan (Septiani & Isabela, 2022). Secara matematis, persamaan pembobotan TF-IDF dirumuskan sebagai berikut :

$$TF-IDF(t, d) = TF(t, d) \times IDF(t) \quad (2.3)$$

keterangan :

- $TF - IDF(t, d)$: Bobot akhir *term t* pada dokumen *d*.
- $TF(t, d)$: Nilai frekuensi kemunculan *term t* pada dokumen *d*.
- $IDF(t)$: Nilai frekuensi dokumen terbalik untuk *term t*

Berdasarkan persamaan 2.3, nilai TF-IDF yang dihasilkan mencerminkan tingkat kepentingan relatif suatu term dalam dokumen terhadap keseluruhan korpus. Dengan demikian, pembobotan TF-IDF menjadi tahapan penting dalam proses transformasi data teks ke dalam bentuk numerik untuk mendukung analisis selanjutnya.

2.5 *Random Forest*

Random Forest merupakan algoritma *ensemble learning* berbasis pohon keputusan (decision tree) yang digunakan untuk tugas klasifikasi maupun regresi (Morama et al., 2022). Metode ini bekerja dengan membangun sejumlah pohon keputusan secara independen, kemudian menggabungkan hasil prediksi masing-masing pohon untuk menghasilkan keputusan akhir (Morama et al., 2022). Setiap pohon dilatih secara independen sehingga menghasilkan struktur keputusan yang berbeda satu sama lain. Keberagaman model ini menjadi kunci utama dalam

meningkatkan kemampuan generalisasi. Semakin rendah korelasi antar pohon, semakin baik performa agregasi yang dihasilkan (Morama et al., 2022).

Proses pembentukan setiap pohon diawali dengan teknik *bootstrap sampling*. *Bootstrap sampling* merupakan teknik pengambilan sampel secara acak dari dataset asli dengan prinsip pengembalian (*with replacement*) (Baskoro et al., 2021). Teknik ini merupakan komponen utama dalam strategi *Bagging (Bootstrap Aggregating)* yang bertujuan untuk membangun keberagaman model dalam *ensemble*. Dari dataset pelatihan D dengan jumlah data sebanyak N , dilakukan proses pengambilan sampel secara repetitif untuk menghasilkan T buah subset data. Setiap subset data hasil *bootstrap* memiliki ukuran yang sama dengan dataset asli (N), namun dengan komposisi data yang berbeda. Karena dilakukan dengan pengembalian, terdapat kemungkinan satu data terpilih lebih dari satu kali dalam satu subset yang sama, sementara data lainnya mungkin tidak terpilih sama sekali (*Out-of-Bag*) (Baskoro et al., 2021). Keberagaman komposisi antar subset inilah yang menjadi dasar pembentukan banyak pohon keputusan dalam *Random Forest* guna meningkatkan stabilitas dan akurasi model. Secara formal, subset data yang dihasilkan dari dataset pelatihan awal dinotasikan sebagai berikut :

$$D_T = \{x_1, x_2, \dots, x_N\} \quad (2.4)$$

keterangan :

- D_T : Subset data (kumpulan ulasan) untuk melatih pohon ke- T .
- x_i : Data ulasan ke- i yang terpilih masuk ke dalam subset melalui proses acak.
- N : Jumlah total baris ulasan dalam satu subset dengan jumlah data sama dengan data set awal.
- T : Target jumlah pohon yang akan dibuat.

Selain variasi data, *Random Forest* juga menciptakan variasi fitur melalui mekanisme *random feature selection*. *Random Feature Selection* adalah mekanisme pemilihan fitur secara acak yang dilakukan pada setiap *node* selama proses pembangunan pohon keputusan (Morama et al., 2022). Berbeda dengan pohon keputusan standar yang mencari pembagi (*split*) terbaik dari seluruh fitur yang tersedia, *Random Forest* hanya mempertimbangkan sebagian kecil fitur acak. Strategi ini bertujuan untuk menekan korelasi antar pohon (*de-correlation*), sehingga jika terdapat satu fitur yang sangat dominan, fitur tersebut tidak akan selalu muncul di setiap pohon dalam *forest* (Morama et al., 2022). Jika p adalah total fitur dari TF-IDF maka syarat pemilihan fitur untuk pemilihan secara acak (m) adalah $m \leq p$. Dalam tugas klasifikasi, nilai standar yang umum digunakan untuk menentukan jumlah fitur acak (m) adalah akar kuadrat dari total fitur yang tersedia, yang dirumuskan sebagai berikut:

$$m = \sqrt{p} \quad (2.5)$$

keterangan :

- m : Jumlah fitur yang dipilih secara acak untuk dipertimbangkan pada setiap *node*.
- p : Total seluruh fitur (dimensi kolom) yang dihasilkan dari tahap pembobotan TF-IDF.

Setelah fitur acak terpilih, setiap pohon keputusan dalam model *Random Forest* dibangun menggunakan subset data yang diperoleh melalui proses *bootstrap sampling*. Pendekatan ini memungkinkan setiap pohon memiliki variasi distribusi data sehingga meningkatkan keberagaman model. Pada proses pembentukan struktur pohon, pemisahan *node* dilakukan dengan menggunakan kriteria Gini

Impurity untuk menentukan fitur terbaik dalam membagi data (Larasati et al., 2022). Secara matematis, nilai *Gini Impurity* dirumuskan sebagai berikut:

$$Gini(t) = 1 - \sum_{i=1}^C p_i^2 \quad (2.6)$$

keterangan:

- p_i = proporsi kelas ke-i pada node t
- C = jumlah kelas

Dengan kombinasi *bootstrap sampling* dan perhitungan *Gini Impurity*, setiap pohon dalam *Random Forest* dapat terbentuk secara beragam namun tetap mempertahankan kemampuan klasifikasi yang akurat. Dalam konteks klasifikasi, setiap pohon yang telah terbentuk melalui proses *training* akan berfungsi sebagai *base learner* yang menghasilkan satu prediksi kelas atau output hipotesis untuk data masukan tertentu. Prediksi akhir dari model *ensemble* diperoleh melalui mekanisme *Majority Voting*, yaitu kelas dengan frekuensi tertinggi dari seluruh output pohon akan dipilih sebagai keputusan akhir model (Morama et al., 2022). Dalam mekanisme ini, jika terdapat T pohon dan masing-masing pohon menghasilkan prediksi $h_t(x)$, maka seluruh hasil prediksi tersebut akan dikumpulkan. Prediksi akhir ditentukan berdasarkan kelas yang memiliki jumlah kemunculan terbanyak di antara seluruh pohon. Secara matematis, keputusan akhir dapat dinyatakan sebagai berikut :

$$\hat{y} = \text{mode}\{h_1(x).h_2(x)....h_T(x)\} \quad (2.7)$$

keterangan :

- $\hat{y}$: Hasil prediksi akhir kelas sentimen (output model Random Forest).
- $h_1(x)$: Hasil prediksi atau output dari pohon keputusan ke-1 untuk input
- $h_T(x)$: Hasil prediksi atau output dari pohon keputusan ke-T untuk input

Pendekatan agregatif ini memberikan stabilitas prediksi yang lebih tinggi dibandingkan penggunaan satu pohon keputusan tunggal. Dengan menggabungkan banyak hipotesis, *Random Forest* mampu mereduksi varians secara signifikan, sehingga model menjadi lebih tangguh terhadap *noise* dan memiliki kemampuan generalisasi yang lebih baik pada data ulasan yang belum pernah dilihat sebelumnya (Morama et al., 2022).

2.6 Evaluasi Performa Klasifikasi

Evaluasi performa merupakan tahap krusial untuk mengukur sejauh mana model klasifikasi mampu memprediksi kelas sentimen secara akurat. Dalam penelitian ini, efektivitas algoritma *Random Forest* diukur menggunakan instrumen pengujian standar yang membandingkan hasil prediksi model dengan label aktual pada dataset.

2.6.1 *Confusion Matrix*

Confusion Matrix adalah tabel kontingensi yang digunakan untuk memvisualisasikan performa algoritma klasifikasi (Morama et al., 2022). Tabel ini menyajikan informasi mengenai perbandingan antara label aktual (kondisi sebenarnya) dengan label prediksi yang dihasilkan oleh model. Dengan menggunakan *Confusion Matrix*, peneliti dapat mengidentifikasi jenis kesalahan yang dilakukan oleh model, baik itu kesalahan dalam memprediksi kelas positif maupun negatif (Morama et al., 2022).

Dalam klasifikasi biner (Positif dan Negatif), *Confusion Matrix* terdiri dari empat komponen utama (Supriyono et al., 2026):

- a. *True Positive* (TP): Jumlah data positif yang diprediksi dengan benar sebagai positif.
- b. *True Negative* (TN): Jumlah data negatif yang diprediksi dengan benar sebagai negatif.
- c. *False Positive* (FP): Jumlah data negatif yang salah diprediksi sebagai positif (*Type I Error*).
- d. *False Negative* (FN): Jumlah data positif yang salah diprediksi sebagai negatif (*Type II Error*).

2.6.2 Metrik Pengukuran Kinerja

Berdasarkan nilai-nilai yang terdapat pada *Confusion Matrix*, dapat dihitung beberapa metrik utama untuk menentukan kualitas model secara kuantitatif. Metrik yang umum digunakan meliputi *Accuracy*, *Precision*, *Recall*, dan *F1-Score* (Supriyono et al., 2026).

a. *Accuracy* (Akurasi)

Akurasi merupakan metrik yang paling intuitif, yang merepresentasikan derajat kedekatan antara hasil prediksi model dengan nilai aktualnya (Larasati et al., 2022). Metrik ini memberikan gambaran umum mengenai sejauh mana model mampu mengklasifikasikan data sentimen (baik positif maupun negatif) secara tepat dari keseluruhan *dataset*. Secara matematis, akurasi dirumuskan sebagai berikut :

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

Berdasarkan persamaan 2.8, nilai akurasi dapat memberikan hasil yang menyesatkan (*accuracy paradox*) jika dataset memiliki ketimpangan jumlah kelas yang signifikan, karena model cenderung berpihak pada kelas mayoritas.

b. *Precision* (Presisi)

Presisi atau disebut juga *Positive Predictive Value* (PPV) adalah metrik yang berfokus pada kualitas prediksi positif model (Larasati et al., 2022). Presisi mengukur tingkat ketepatan model dengan cara menghitung rasio ulasan yang benar-benar positif dibandingkan dengan seluruh ulasan yang "diklaim" positif oleh model. Secara matematis, presisi dirumuskan sebagai berikut :

$$Precision = \frac{TP}{TP + FP} \quad (2.9)$$

Berdasarkan persamaan 2.9, nilai presisi yang tinggi menunjukkan bahwa model memiliki tingkat kesalahan *False Positive* (FP) yang rendah. Dalam konteks ulasan wisata, presisi yang baik memastikan bahwa ulasan yang diprediksi "Positif" memang benar-benar berisi pujian, bukan ulasan negatif yang salah terdeteksi.

c. *Recall* (Sensitivitas)

Recall atau *True Positive Rate* (TPR) mengukur kemampuan model dalam mengidentifikasi seluruh anggota dari kelas positif (Larasati et al., 2022). Berbeda dengan presisi, *recall* menitikberatkan pada seberapa banyak ulasan positif yang berhasil "ditangkap" kembali oleh model dari total ulasan positif yang sebenarnya ada. Secara matematis, *recall* dirumuskan sebagai berikut :

$$Recall = \frac{TP}{TP + FN} \quad (2.10)$$

Berdasarkan persamaan 2.10, nilai *recall* yang rendah menunjukkan adanya banyak kesalahan *False Negative* (FN), yang berarti banyak ulasan sentimen positif yang justru terlewatkan dan diklasifikasikan sebagai negatif oleh model.

d. *F1-Score*

F1-Score merupakan rata-rata harmonik (*harmonic mean*) dari presisi dan *recall* (Larasati et al., 2022). Metrik ini memberikan bobot yang seimbang antara kedua metrik tersebut dan sering digunakan sebagai indikator utama ketika peneliti mencari titik tengah antara optimasi presisi dan *recall* (Supriyono et al., 2024). Secara matematis, *f1-score* dirumuskan sebagai berikut:

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.11)$$

Berdasarkan persamaan 2.11, jika salah satu dari presisi atau *recall* bernilai sangat rendah, maka nilai *F1-Score* akan turun secara drastis. Hal ini menjadikannya parameter yang sangat reliabel untuk mengukur stabilitas model dalam menangani data sentimen yang kompleks.

2.7 Penelitian Terkait

Kajian terhadap penelitian terdahulu dilakukan untuk memahami perkembangan pendekatan *Aspect-Based Sentiment Analysis* (ABSA), khususnya yang memanfaatkan kombinasi teknik representasi fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) dan algoritma klasifikasi *machine learning*. Analisis ini bertujuan untuk mengidentifikasi posisi penelitian yang dilakukan, sekaligus menemukan *research gap* yang belum banyak dieksplorasi dalam domain

pariwisata terpadu. Data ringkasan mengenai perbandingan antara penelitian terdahulu dengan rencana penelitian ini disajikan dalam Tabel 2.1 berikut:

Tabel 2.1 Penelitian Terdahulu

No	Peneliti Dan Judul	Metode	Keterbatasan Penelitian	Research Gap
1	Firdaus Dkk. (2025), “Analisis Sentimen Berbasis Aspek Ulasan Pengguna Aplikasi Alfagift Menggunakan Metode Random Forest Dan Pemodelan Topik Latent Dirichlet Allocation”	ABSA ulasan aplikasi belanja menggunakan <i>Random Forest</i> , LDA, dan TF-IDF.	Berfokus pada domain <i>e-commerce</i> ; karakteristik ulasan cenderung teknis terkait aplikasi.	Menerapkan <i>Random Forest</i> pada domain pariwisata untuk menguji ketangguhan model pada teks subjektif.
2	Putra & Ratnawati (2025), “Analisis Sentimen Berbasis Aspek Pada Aplikasi Mobile Menggunakan Naïve Bayes Berdasarkan Ulasan Pengguna Playstore (Studi Kasus : Jconnect Mobile)”	ABSA aplikasi <i>mobile banking</i> menggunakan <i>Naïve Bayes</i> , TF-IDF, dan RCA.	Menggunakan algoritma tunggal (<i>single learner</i>); fokus pada identifikasi <i>bug</i> sistem.	Menguji performa <i>Random Forest</i> (ensemble) untuk klasifikasi aspek yang lebih beragam berdasarkan konsep 3A pariwisata.
3	M. Syamsul Hadi Dkk. (2025),	Analisis sentimen	Klasifikasi sentimen	Mengembangkan analisis berbasis

No	Peneliti Dan Judul	Metode	Keterbatasan Penelitian	Research Gap
	“Analisis Sentimen Wisata Air Terjun Di Kabupaten Lombok Tengah Menggunakan Metode Support Vector Machine (SVM)	ulasan Google Maps menggunakan <i>Support Vector Machine</i> (SVM).	dilakukan secara global; belum membedah ke tingkat aspek (ABSA).	aspek untuk memberikan rincian evaluasi kualitas destinasi.
4	Parasati Dkk. (2020), “Analisis Sentimen Berbasis Aspek Pada Ulasan Pelanggan Restoran Bakso President Malang Dengan Metode Naïve Bayes Classifier”	ABSA domain kuliner di Malang menggunakan <i>Naïve Bayes</i> dan TF-IDF.	Konteks objek wisata tunggal; keterbatasan akurasi pada algoritma probabilitas sederhana.	Mengembangkan analisis pada destinasi <i>multi-site</i> (Jatim Park Group) dengan volume data lebih masif.

Berdasarkan perbandingan pada Tabel 2.1, penelitian pertama oleh Firdaus dkk. (2025) memberikan landasan teknis mengenai efektivitas algoritma *Random Forest* dalam tugas ABSA. Studi tersebut menunjukkan bahwa kombinasi TF-IDF dan *Random Forest* mampu menangani *dataset* ulasan yang kompleks dengan akurasi mencapai 89%. Relevansi utama penelitian tersebut dengan penelitian ini terletak pada penggunaan model klasifikasi yang sama, namun terdapat perbedaan fundamental pada karakteristik domain data yang digunakan.

Penelitian berikutnya oleh Putra dan Ratnawati (2025) menerapkan ABSA pada aplikasi perbankan menggunakan *Naïve Bayes*. Meskipun menghasilkan akurasi yang tinggi sebesar 93,1%, penelitian tersebut mencatat bahwa analisis

sentimen lebih diarahkan pada perbaikan teknis sistem melalui *Root Cause Analysis* (RCA). Hal ini memberikan perspektif bagi peneliti untuk mengarahkan hasil klasifikasi ulasan Jatim Park Group menjadi rekomendasi strategis bagi manajemen manajemen destinasi, bukan sekadar angka akurasi.

Sementara itu, penelitian oleh M. Syamsul Hadi dkk. (2025) memperkuat argumen penggunaan ulasan *Google Maps* sebagai sumber data yang valid untuk menilai kualitas tempat wisata. Studi tersebut menggunakan SVM untuk memetakan kepuasan pengunjung di Lombok Tengah. Namun, sebagaimana terlihat pada kolom keterbatasan Tabel 2.1, studi tersebut hanya melakukan klasifikasi sentimen secara umum. Penelitian ini berupaya mengisi celah tersebut dengan membedah ulasan ke dalam aspek spesifik *Attraction*, *Amenity*, dan *Accessibility*.

Penelitian oleh Parasati dkk. (2020) memiliki kemiripan konteks lokasi di Malang Raya dan industri pariwisata. Dengan menggunakan *Naïve Bayes*, penelitian tersebut membedah aspek ulasan pada satu titik lokasi restoran. Peneliti menggunakan studi ini sebagai pembandingan metode, di mana *Random Forest* yang digunakan dalam penelitian ini diproyeksikan memiliki performa yang lebih stabil dalam menangani variasi bahasa pada ulasan wisata yang berjumlah ribuan dan memiliki tingkat subjektivitas tinggi.

Berdasarkan tinjauan di atas, ditemukan bahwa penelitian ABSA menggunakan *Random Forest* mayoritas masih diimplementasikan pada domain aplikasi digital. Terdapat kekosongan literatur pada penerapan metode ini untuk mengevaluasi destinasi wisata tematik berskala besar seperti Jatim Park Group yang memiliki volume ulasan masif dan karakteristik data yang lebih "berisik" (*noisy*). Dengan mengintegrasikan algoritma *Random Forest* dan kerangka aspek 3A pariwisata, penelitian ini diharapkan dapat memberikan kontribusi metodologis dalam memperluas domain penerapan ABSA serta memberikan hasil evaluasi yang lebih akurat bagi industri pariwisata.

BAB III

DESAIN PENELITIAN

3.1 Jenis Penelitian

Penelitian ini merupakan penelitian kuantitatif dengan metode eksperimen. Pendekatan kuantitatif digunakan karena proses analisis dilakukan melalui pengolahan data teks yang ditransformasikan ke dalam bentuk numerik, kemudian dianalisis menggunakan model klasifikasi berbasis *machine learning*.

Metode eksperimen diterapkan untuk menguji performa algoritma dalam mengklasifikasikan sentimen berbasis aspek pada ulasan wisata, pengujian dilakukan dengan membangun model klasifikasi dan mengukur kinerjanya menggunakan metrik evaluasi statistik. Penelitian ini berfokus pada evaluasi performa model dan tidak mencakup perancangan sistem aplikasi maupun pembahasan strategi manajerial pengelolaan objek wisata.

Pendekatan ABSA memungkinkan proses analisis sentimen dilakukan secara lebih terperinci dengan memetakan opini pengguna ke dalam kategori aspek tertentu. Dengan skema ini, setiap ulasan tidak hanya diperlakukan sebagai satu kesatuan sentimen umum, tetapi dianalisis berdasarkan aspek spesifik yang relevan.

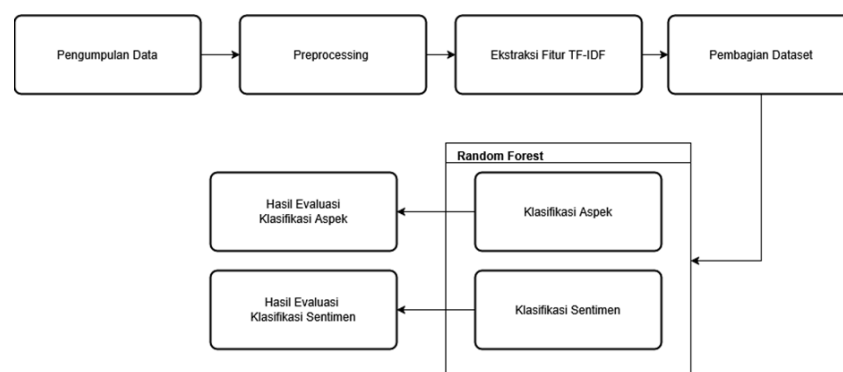
Setelah proses identifikasi dan anotasi aspek dilakukan, tahap berikutnya adalah klasifikasi sentimen pada setiap aspek yang telah ditentukan. Setiap data yang telah diberi label aspek kemudian diklasifikasikan ke dalam polaritas sentimen. Proses klasifikasi dilakukan menggunakan algoritma *Random Forest* sebagai model pembelajaran terawasi. Pemilihan algoritma ini didasarkan pada

kemampuannya dalam menangani data berdimensi tinggi serta mengurangi risiko *overfitting* melalui mekanisme *ensemble learning*. Dengan demikian, pendekatan penelitian ini mengintegrasikan identifikasi aspek dan klasifikasi sentimen dalam satu alur analisis yang sistematis.

3.2 Perancangan Alur Penelitian

Perancangan alur penelitian ini disusun untuk menggambarkan tahapan sistematis dalam pembangunan sistem klasifikasi *multi-label* menggunakan satu model pembelajaran mesin. Alur penelitian dirancang secara terstruktur mulai dari tahap persiapan data hingga proses evaluasi model, sehingga setiap tahapan memiliki keterkaitan yang jelas dan mendukung pencapaian tujuan penelitian.

Penelitian ini menggunakan pendekatan klasifikasi *multi-output multi-class*. Melalui pendekatan ini, setiap data menghasilkan beberapa luaran klasifikasi yang saling independen. Setiap luaran merepresentasikan satu aspek yang memiliki empat kemungkinan kelas sentimen, yaitu Positif, Negatif, Netral, dan *None*. Oleh karena itu, sistem dirancang menggunakan satu model tunggal yang mampu memprediksi seluruh aspek secara keseluruhan dalam satu proses. Alur penelitian digambarkan pada Gambar 3.1.



Gambar 3.1 Perancangan Alur Penelitian

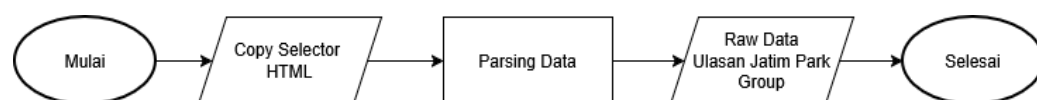
Berdasarkan alur pada Gambar 3.1, setiap tahapan dilakukan secara berurutan dan saling terintegrasi, mulai dari pengolahan data hingga evaluasi performa model. Perancangan alur ini memastikan bahwa proses penelitian berjalan secara sistematis, terkontrol, dan sesuai dengan metodologi yang telah ditetapkan untuk menghasilkan model klasifikasi *multi-label* yang optimal.

3.3 Pengumpulan Data

Setelah seluruh data ulasan dari platform *Google Maps* dan *TripAdvisor* berhasil dihimpun, langkah selanjutnya adalah melakukan pengorganisasian *dataset* untuk menjamin integritas informasi yang akan diolah. Proses ini melibatkan sinkronisasi format antar sumber data yang berbeda guna menciptakan satu kesatuan basis data yang koheren. Dengan tersusunnya *dataset* mentah ini, peneliti memiliki landasan empiris yang kuat untuk melakukan tahapan seleksi dan pembersihan, sehingga karakteristik data yang dihasilkan benar-benar merepresentasikan opini pengunjung terhadap unit-unit rekreasi di Jatim Park Group secara akurat.

3.3.1 Teknik Pengumpulan Data

Pengumpulan data dilakukan menggunakan metode ekstraksi berbasis struktur halaman web (*web scraping*) dengan pendekatan seleksi elemen HTML (*HTML selector*).



Gambar 3.2 Proses Pengumpulan Data

Tahapan pengumpulan data berdasarkan Gambar 3.2 meliputi:

1. Mengidentifikasi elemen HTML yang memuat teks ulasan pada masing-masing platform menggunakan teknik selektor HTML.
2. Kemudian *Parsing* Data dari teks HTML menggunakan skrip berbasis *Python* untuk mengumpulkan dan menyusun data secara terstruktur.
3. Menyimpan *Raw Data* hasil ekstraksi dalam format *spreadsheet* (Excel) sebagai *dataset* awal penelitian.

Setelah data dari kedua platform terkumpul, dilakukan proses penggabungan *dataset* dan pembersihan data awal sebelum memasuki tahap anotasi dan *preprocessing*. Proses ini dilakukan dengan memanfaatkan data yang bersifat publik serta tidak menyertakan informasi pribadi pengguna, sehingga tetap memperhatikan prinsip etika penelitian.

3.3.2 Karakteristik dan Riwayat Dataset

Dataset dalam penelitian ini melalui beberapa tahapan seleksi guna menjamin kualitas data yang akan diproses oleh model. Riwayat *dataset* mencakup sumber platform, sebaran objek wisata, serta periode pengambilan data sebagai berikut:

1. Sumber Platform dan Objek Wisata

Data diekstraksi dari dua platform ulasan daring utama, yaitu *Google Maps* dan *TripAdvisor*. Objek wisata yang menjadi fokus penelitian adalah unit-unit rekreasi di bawah naungan Jatim Park Group yang berlokasi di Kota Batu, Jawa Timur. Unit rekreasi tersebut meliputi:

- Jatim Park 1
- Jatim Park 2
- Jatim Park 3
- Museum Angkut

2. Periode Pengambilan Data

Rentang waktu pengambilan data ulasan yang diambil mencakup periode waktu dari 2020 hingga 2025.

3. Karakteristik dan Jumlah Dataset

Riwayat penyusutan *dataset* terjadi secara bertahap seiring dengan proses seleksi kualitas data. Hasil ekstraksi diperoleh 4.777 ulasan mentah dari kedua platform yang disajikan pada tabel 3.1 berikut :

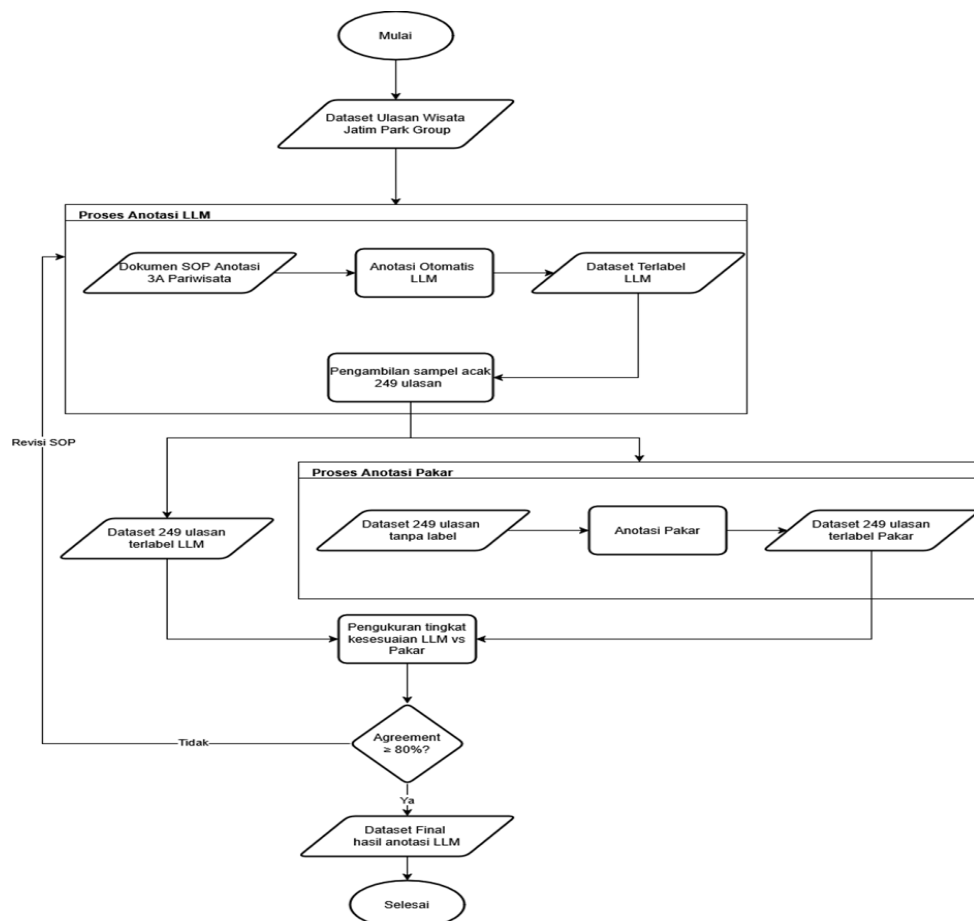
Tabel 3.1 Rincian Data dari Platform

Objek Wisata	Platform	
	Gmaps	TripAdvisor
Jatim Park 1	757	323
Jatim Park 2	1871	487
Jatim Park 3	525	209
Museum Angkut	605	0
Total	3758	1019

Kemudian pembersihan data yang meliputi penghapusan ulasan duplikat serta karakter yang tidak terbaca (*corrupted text*). Proses ini menyusutkan dataset awal pada Tabel 3.1 dari 4.777 menjadi 4.195 ulasan, yang kemudian dilanjutkan ke tahap anotasi.

3.3.3 Anotasi Data

Tahap anotasi data dilakukan untuk memberikan label sentimen terhadap setiap ulasan berdasarkan pendekatan 3A dalam domain pariwisata. *Dataset* yang digunakan dalam penelitian ini terdiri dari 4.195 ulasan berbentuk teks yang telah melalui proses seleksi awal. Anotasi bertujuan untuk mengidentifikasi polaritas sentimen pada masing-masing aspek sehingga data siap digunakan dalam proses pelatihan model klasifikasi, proses ini menjadi tahap krusial karena kualitas anotasi akan memengaruhi performa model yang dibangun. Alur proses anotasi data dapat dilihat pada Gambar 3.3 berikut :



Gambar 3.3 Proses Anotasi Data

Berdasarkan Gambar 3.3 mengilustrasikan alur proses anotasi dataset ulasan wisata Jatim Park Group secara keseluruhan. Proses dimulai dari input dataset ulasan wisata Jatim Park Group yang kemudian masuk ke dalam dua tahap proses utama secara bertahap.

Tahap pertama adalah Proses Anotasi LLM. Pada tahap ini, Dokumen SOP Anotasi 3A Pariwisata digunakan sebagai input panduan bagi model LLM untuk menjalankan proses anotasi otomatis terhadap seluruh dataset. Setiap ulasan dianotasi terhadap tiga aspek yaitu atraksi, amenitas, dan aksesibilitas. Untuk setiap aspek, ditetapkan salah satu dari empat kelas sentimen, yaitu positif, negatif, netral, atau none apabila ulasan tidak menyinggung aspek yang bersangkutan. Hasil dari proses ini adalah dataset terlabel LLM yang mencakup keseluruhan data. Dari dataset tersebut kemudian dilakukan pengambilan sampel acak sebanyak 249 ulasan atau sekitar 5% dari keseluruhan dataset sebagai subset untuk keperluan validasi.

Tahap kedua adalah Proses Anotasi Pakar. Sebanyak 249 ulasan tanpa label diserahkan kepada pakar untuk dianotasi secara manual dan independen menggunakan pedoman anotasi yang sama. Hasil dari proses ini adalah Dataset 249 ulasan terlabel pakar.

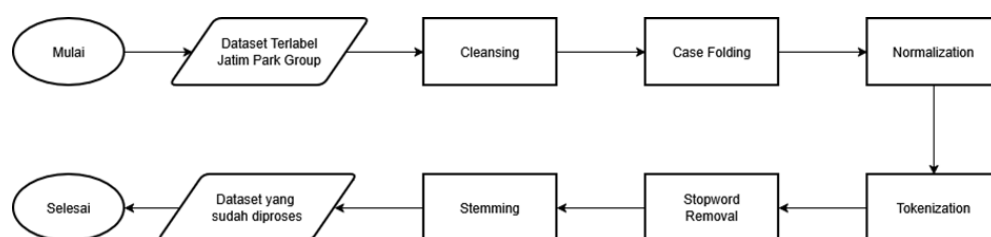
Setelah kedua proses selesai, dua output dibandingkan secara bersamaan, yaitu dataset 249 ulasan terlabel LLM dan dataset 249 ulasan terlabel Pakar, melalui proses pengukuran tingkat kesesuaian LLM vs Pakar. Hasil pengukuran kemudian dievaluasi melalui nilai kesesuaian dengan ambang batas kesesuaian $\geq 80\%$. Apabila nilai kesesuaian tidak mencapai ambang batas tersebut, proses

dikembalikan ke tahap anotasi LLM untuk dilakukan revisi SOP. Apabila nilai kesesuaian memenuhi ambang batas, maka dihasilkan dataset Final hasil anotasi LLM yang selanjutnya siap digunakan pada tahap pelatihan dan pengujian model klasifikasi.

Hasil pengukuran tingkat kesesuaian antara anotasi LLM dan anotasi Pakar menunjukkan nilai rata-rata sebesar 82%. Proses validasi ini dilakukan oleh Yoga Yolanda, M.Pd., dosen bahasa Indonesia dari UIN Maulana Malik Ibrahim Malang yang melakukan anotasi secara independen tanpa mengetahui hasil label LLM sebelumnya. Nilai kesesuaian tersebut melampaui ambang batas yang ditetapkan sebesar 80%, sehingga dataset final hasil anotasi LLM ditetapkan sebagai dataset yang digunakan pada tahap pelatihan dan pengujian model klasifikasi.

3.4 *Preprocessing*

Tahapan *preprocessing* merupakan proses krusial dalam *text mining* yang bertujuan mentransformasi data teks mentah yang bersifat tidak terstruktur menjadi format data terstruktur. Mengingat data mentah tidak dapat diolah secara langsung oleh algoritma, rangkaian pemrosesan awal ini diperlukan agar data memiliki kualitas yang memadai untuk dianalisis lebih lanjut (Deviyanto & Wahyudi, 2018). Alur proses preprocessing data pada penelitian ini dapat dilihat pada Gambar 3.4 berikut :



Gambar 3.4 Alur *PreProcessing* Data

Berdasarkan Gambar 3.4 penjelasan langkah-langkah text preprocessing yang digunakan dalam penelitian ini adalah sebagai berikut :

a. Cleansing

Tahap cleansing dalam penelitian ini bertujuan untuk membersihkan teks dari elemen-elemen yang tidak diperlukan dalam analisis, seperti tautan, angka, tanda baca, simbol, dan spasi berlebih. Proses ini dilakukan menggunakan *regular expression* (regex), yaitu pola teks tertentu yang digunakan untuk mendeteksi dan menghapus karakter yang tidak diinginkan. Penghapusan dilakukan dalam dua langkah: pertama menghapus tautan URL menggunakan pola `http\S+`, kemudian menghapus angka dan simbol menggunakan pola `[0-9]+[^\w\s]`. Karakter yang cocok dengan pola tersebut diganti dengan spasi kosong, lalu spasi berlebih di awal dan akhir teks dibersihkan menggunakan fungsi `.strip()`. Implementasi *Cleansing* dapat dilihat pada Tabel 3.3.

Tabel 3.2 Proses *Cleansing*

Sebelum Cleansing	Setelah Cleansing
Seru...wahananya gk terlalu banyak,tp cukup utk di nikmati dr pagi smpe sore... Gak capek jg krn gk trllu banyak jalan 🍷 🍷	Seru wahananya gk terlalu banyak tp cukup utk di nikmati dr pagi smpe sore Gak capek jg krn gk trllu banyak jalan

Berdasarkan Tabel 3.3, terlihat bahwa tanda baca seperti tanda titik, koma, dan elemen emotikon 🍷 🍷 berhasil dihilangkan dari teks. Spasi berlebih akibat

penghapusan karakter tersebut juga telah dirapikan, sehingga teks menjadi lebih bersih dan siap untuk diproses pada tahap selanjutnya.

b. Case Folding

Tahap *case folding* dalam penelitian ini bertujuan untuk mengubah seluruh huruf dalam teks menjadi huruf kecil. Hal ini penting agar kata yang sama tidak dianggap berbeda hanya karena perbedaan huruf kapital, misalnya "Bagus" dan "bagus" akan diperlakukan sebagai kata yang sama. Secara teknis, setiap karakter dalam teks diperiksa berdasarkan nilai ASCII-nya. Huruf kapital memiliki nilai ASCII 65 hingga 90 (A sampai Z), dan dikonversi menjadi huruf kecil dengan menambahkan nilai 32, sehingga menghasilkan nilai ASCII 97 hingga 122 (a sampai z). Implementasi *case folding* dapat dilihat pada Tabel 3.2.

Tabel 3.3 Proses *Case Folding*

Sebelum Case Folding	Setelah Case Folding
Seru wahananya gk terlalu banyak tp cukup utk di nikmati dr pagi smpe sore Gak capek jg krn gk trllu banyak jalan	seru wahananya gk terlalu banyak tp cukup utk di nikmati dr pagi smpe sore gak capek jg krn gk trllu banyak jalan

Berdasarkan Tabel 3.2, terlihat bahwa seluruh huruf kapital pada teks, seperti "Seru" dan "Gak", telah diubah menjadi huruf kecil "seru" dan "gak". Dengan demikian, kata yang sama tidak lagi dibedakan berdasarkan penulisan huruf besar atau kecil.

c. Normalization

Tahap normalisasi dalam penelitian ini bertujuan untuk mengubah kata-kata slang atau tidak baku menjadi bentuk formalnya. Dalam proses ini diperlukan sebuah kamus slang sebagai acuan penggantian kata. Kamus yang digunakan adalah

Colloquial Indonesian Lexicon yang dikembangkan oleh Salsabila et al. (2018) tersedia secara terbuka di GitHub dengan lisensi MIT. Cara kerjanya adalah setiap kata dalam teks dicocokkan ke dalam kamus. Jika kata tersebut ditemukan maka diganti dengan bentuk formalnya, jika tidak ditemukan maka kata dibiarkan seperti semula. Implementasi *Normalization* dapat dilihat pada Tabel 3.4.

Tabel 3.4 Proses *Normalization*

Sebelum Normalization	Setelah Normalization
seru wahananya gk terlalu banyak tp cukup utk di nikmati dr pagi smpe sore gak capek jg krn gk trllu banyak jalan	seru wahananya tidak terlalu banyak tapi cukup untuk di nikmati dari pagi sampai sore tidak capek juga karena tidak terlalu banyak jalan

Berdasarkan Tabel 3.4, terlihat bahwa kata-kata slang seperti "gk", "tp", "utk", "dr", "smpe", "jg", "krn", dan "trllu" telah diubah menjadi bentuk formalnya, yaitu "tidak", "tapi", "untuk", "dari", "sampai", "juga", "karena", dan "terlalu". Perubahan ini membuat representasi kata menjadi lebih konsisten dan mengurangi variasi penulisan untuk makna yang sama.

d. *Tokenization*

Tahap *tokenization* dalam penelitian ini bertujuan untuk memecah teks menjadi satuan kata-kata individual yang disebut token. Proses ini bekerja dengan cara memindai teks dan memotongnya setiap kali menemukan spasi sebagai pemisah antar kata, sehingga menghasilkan daftar kata. Misalnya kalimat "wahana sangat seru" akan dipecah menjadi ["wahana", "sangat", "seru"]. Hasil tokenisasi ini kemudian digunakan sebagai *input* pada tahap *stopword removal* dan *stemming*. Implementasi *Tokenization* dapat dilihat pada Tabel 3.5.

Tabel 3.5 Proses *Tokenization*

Sebelum Tokenization	Setelah Tokenization
seru wahananya tidak terlalu banyak tapi cukup untuk di nikmati dari pagi sampai sore tidak capek juga karena tidak terlalu banyak jalan	["seru", "wahananya", "tidak", "terlalu", "banyak", "tapi", "cukup", "untuk", "di", "nikmati", "dari", "pagi", "sampai", "sore", "tidak", "capek", "juga", "karena", "tidak", "terlalu", "banyak", "jalan"]

Berdasarkan Tabel 3.5, terlihat bahwa kalimat hasil normalisasi telah dipecah menjadi 22 token kata individual dalam bentuk list. Setiap kata kini menjadi satuan terpisah yang dapat diproses lebih lanjut pada tahap *stopword removal*.

e. *Stopword Removal*

Tahap *stopword removal* dalam penelitian ini bertujuan untuk menghapus kata-kata yang tidak memiliki kontribusi makna dalam analisis sentimen, seperti kata "yang", "di", "dan", dan sejenisnya. Proses ini bekerja dengan cara memeriksa setiap token satu per satu: apabila token terdapat dalam daftar *stopword* Bahasa Indonesia dari library Sastrawi maka token tersebut dihapus, apabila tidak maka token dipertahankan. Selain itu, token yang hanya terdiri dari satu karakter juga langsung dihapus karena dianggap tidak bermakna. Implementasi *Stopword Removal* dapat dilihat pada Tabel 3.6.

Tabel 3.6 Proses *Stopword Removal*

Sebelum Stopword Removal	Setelah Stopword Removal
["seru", "wahananya", "tidak", "terlalu", "banyak", "tapi", "cukup",	["seru", "wahananya", "terlalu", "banyak", "cukup", "nikmati", "pagi",

Sebelum Stopword Removal	Setelah Stopword Removal
"untuk", "di", "nikmati", "dari", "pagi", "sampai", "sore", "tidak", "capek", "juga", "karena", "tidak", "terlalu", "banyak", "jalan"]	"sore", "capek", "terlalu", "banyak", "jalan"]

Berdasarkan Tabel 3.6, terlihat bahwa kata-kata seperti "tidak", "tapi", "untuk", "di", "dari", "juga", dan "karena" telah dihapus karena termasuk dalam daftar *stopword* maupun token dengan satu karakter. Jumlah token berkurang dari 22 menjadi 12, sehingga teks yang tersisa hanya berisi kata-kata yang memiliki makna terhadap analisis sentimen.

f. Stemming

Tahap stemming dalam penelitian ini bertujuan untuk mengubah kata berimbuhan menjadi bentuk dasarnya, sehingga kata seperti "menikmati", "dinikmati", dan "menikmati" semuanya dikenali sebagai satu kata dasar yang sama yaitu "nikmat". Proses ini dijalankan menggunakan library PySastrawi yang menerapkan algoritma Nazief-Adriani. Algoritma ini bekerja dengan cara mendeteksi dan melepas imbuhan secara bertahap: pertama melepas sufiks yaitu kata akhiran seperti "-kan", "-an", "-nya" dan "-i", kemudian melepas prefiks yaitu kata awalan seperti "me-", "di-", "ber-", dan "ter-", lalu memeriksa apakah hasilnya merupakan kata dasar yang valid. Jika belum valid, proses diulangi dengan kombinasi aturan yang berbeda hingga ditemukan bentuk dasar yang tepat. Implementasi *Stemming* dapat dilihat pada Tabel 3.7.

Tabel 3.7 Proses *Stemming*

Sebelum Stemming	Setelah Stemming
["seru", "wahananya", "terlalu", "banyak", "cukup", "nikmati", "pagi", "sore", "capek", "terlalu", "banyak", "jalan"]	["seru", "wahana", "terlalu", "banyak", "cukup", "nikmat", "pagi", "sore", "capek", "terlalu", "banyak", "jalan"]

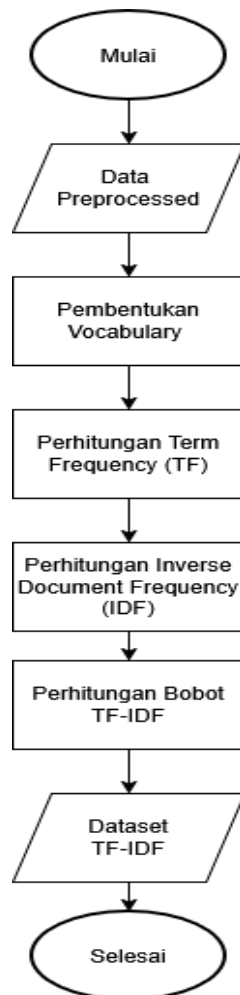
Berdasarkan Tabel 3.7, terlihat bahwa kata "wahananya" diubah menjadi "wahana" dan "nikmati" diubah menjadi "nikmat" setelah imbuhan pada kata tersebut dilepas. Kata-kata lain seperti "seru", "terlalu", "banyak", "cukup", "pagi", "sore", "capek", dan "jalan" tidak mengalami perubahan karena sudah merupakan bentuk dasar.

Data yang telah melewati tahap *preprocessing* ini selanjutnya akan digunakan sebagai input dalam proses ekstraksi fitur menggunakan metode TF-IDF. Tahapan ini berfungsi untuk mengonversi data tekstual yang telah bersih menjadi bentuk matriks numerik, sehingga setiap kata memiliki bobot tertentu yang merepresentasikan tingkat kepentingannya dalam dokumen. Transformasi ini sangat krusial agar data siap diproses oleh algoritma pembelajaran mesin pada tahap analisis selanjutnya.

3.5 Ekstraksi Fitur TF-IDF

Tahap ekstraksi fitur bertujuan untuk mengubah data teks yang telah melalui proses *preprocessing* menjadi representasi numerik dalam bentuk vektor. Pada penelitian ini, metode yang digunakan adalah *Term Frequency-Inverse Document Frequency* (TF-IDF). Berbeda dengan penggunaan pustaka standar, implementasi TF-IDF dalam penelitian ini dilakukan secara mandiri untuk menghitung bobot

setiap kata berdasarkan kemunculannya dalam dokumen dan korpus. Proses kerja dari TF-IDF dapat dilihat pada Gambar 3.5 berikut :



Gambar 3.5 TF-IDF

Implementasi pembobotan TF-IDF diawali dengan menyiapkan dataset yang telah melewati tahap *preprocessing*. Sebagai ilustrasi, digunakan tiga buah dokumen dalam bentuk sekumpulan token seperti yang disajikan pada Tabel 3.8 berikut:

Tabel 3.8 Contoh Dokumen Hasil *Preprocessed*

Dokumen	Review text	Total kata
D1	[indonesia, malang, wahana, budaya]	4
D2	[indonesia, malang, taman, taman, taman, taman, museum, museum, edukasi, kuliner, bagus, bagus]	12
D3	[alam, marah, lingkungan, tanggung, tanggung]	5

Adapun tahapan penerapan TF-IDF akan dilakukan sesuai dengan Gambar 3.5 sebagai berikut :

1. Pembentukan *Vocabulary*

Tahap pertama dalam implementasi TF-IDF adalah pembentukan *vocabulary*. Pada tahap ini, sistem melakukan ekstraksi terhadap seluruh dokumen yang ada untuk mengidentifikasi setiap kata unik yang muncul. Kata-kata unik ini kemudian akan digunakan sebagai fitur (kolom) dalam matriks pembobotan. Berdasarkan tiga dokumen contoh sebelumnya, daftar kata unik yang berhasil dikumpulkan dapat dilihat pada Tabel 3.9 berikut:

Tabel 3.9 Daftar Kosakata (Fitur)

No	Kata (Term)
1	indonesia
2	malang
3	wahana
4	budaya
5	taman
6	museum
7	edukasi
8	kuliner
9	bagus
10	alam
11	marah
12	lingkungan
13	tanggung

2. Perhitungan *Term Frequency* (TF)

Setelah daftar *vocabulary* terbentuk, langkah berikutnya adalah menghitung nilai *Term Frequency* (TF). Nilai TF ini merepresentasikan frekuensi atau jumlah kemunculan setiap kata unik (*term*) di dalam masing-masing dokumen. Berdasarkan persamaan 2.1, nilai $f(t.d)$ merupakan frekuensi kemunculan kata tertentu, sedangkan $\sum_{t' \in d} f(t'.d)$ merupakan total jumlah seluruh kata dalam dokumen tersebut yang digunakan sebagai penyebut (pembagi). Total jumlah kata pada setiap dokumen (D1, D2, dan D3) dari Tabel 3.8 yang menjadi acuan perhitungan dapat dilihat sebagai berikut:

- D1: 4 kata
- D2: 12 kata
- D3: 5 kata

Hasil perhitungan frekuensi kemunculan kata (TF) disajikan pada Tabel 3.10 sebagai berikut:

Tabel 3.10 Hasil Perhitungan Frekuensi Kata (TF)

No	Kata (Term)	TF		
		D1 (f/4)	D2 (f/12)	D3 (f/5)
1	indonesia	0,25	0,083	0
2	malang	0,25	0,083	0
3	wahana	0,25	0	0
4	budaya	0,25	0	0
5	taman	0	0,333	0
6	museum	0	0,166	0
7	edukasi	0	0,083	0
8	kuliner	0	0,083	0
9	bagus	0	0,166	0
10	alam	0	0	0,2
11	marah	0	0	0,2
12	lingkungan	0	0	0,2
13	tanggung	0	0	0,4

3. Perhitungan *Inverse Document Frequency* (IDF)

Tahap berikutnya adalah perhitungan *Inverse Document Frequency* (IDF). IDF digunakan untuk mengukur tingkat kepentingan suatu kata terhadap keseluruhan koleksi dokumen. Konsep ini didasarkan pada asumsi bahwa kata yang muncul di banyak dokumen memiliki daya pembeda yang lebih rendah (kurang penting), sedangkan kata yang jarang muncul memiliki kepentingan yang lebih tinggi. Berdasarkan persamaan 2.2, perhitungan IDF menggunakan logaritma dari total jumlah dokumen (N) dibagi dengan jumlah dokumen yang mengandung kata tersebut ($df(t)$). Hasil perhitungan IDF untuk setiap kata unik disajikan pada Tabel 3.11 berikut:

Tabel 3.11 Hasil Perhitungan *Inverse Document Frequency* (IDF)

No	Kata (Term)	Jumlah Dokumen (df)	Nilai IDF $(\log \frac{3}{df(t)})$
1	indonesia	2	0,176
2	malang	2	0,176
3	wahana	1	0,477
4	budaya	1	0,477
5	taman	1	0,477
6	museum	1	0,477
7	edukasi	1	0,477
8	kuliner	1	0,477
9	bagus	1	0,477
10	alam	1	0,477
11	marah	1	0,477
12	lingkungan	1	0,477
13	tanggung	1	0,477

4. Perhitungan Bobot TF-IDF

Tahap terakhir adalah perhitungan bobot fitur menggunakan metode TF-IDF. Nilai akhir bobot fitur diperoleh dengan mengalikan nilai *Term Frequency* (TF) dengan nilai *Inverse Document Frequency* (IDF). Perkalian ini menghasilkan skor yang menunjukkan seberapa relevan sebuah kata terhadap dokumen

spesifiknya dalam konteks keseluruhan data (korpus). Hasil perhitungan bobot TF-IDF untuk setiap kata unik pada seluruh dokumen disajikan pada Tabel 3.12 sebagai berikut:

Tabel 3.12 Hasil Perhitungan Bobot TF-IDF

No	Kata (Term)	IDF	D1 (TF × IDF)	D2 (TF × IDF)	D3 (TF × IDF)
1	indonesia	0,176	0,176	0,176	0
2	malang	0,176	0,176	0,176	0
3	wahana	0,477	0,477	0	0
4	budaya	0,477	0,477	0	0
5	taman	0,477	0	1,908	0
6	museum	0,477	0	0,954	0
7	edukasi	0,477	0	0,477	0
8	kuliner	0,477	0	0,477	0
9	bagus	0,477	0	0,954	0
10	alam	0,477	0	0	0,477
11	marah	0,477	0	0	0,477
12	lingkungan	0,477	0	0	0,477
13	tanggung	0,477	0	0	0,954

Setelah seluruh tahapan perhitungan selesai, data kini telah berbentuk matriks numerik di mana setiap kata unik (*vocabulary*) bertindak sebagai fitur. Data ini merupakan versi siap pakai untuk proses pelatihan dan pengujian pada model *Random Forest*. Format dataset akhir yang dihasilkan disajikan pada tabel 3.13 sebagai berikut:

Tabel 3.13 Cuplikan *Dataset* Siap Pakai untuk *Random Forest*

No	Teks Ulasan	indonesia	..	Atraksi	Aksesibilitas	Amenitas
1	[indonesia, malang, wahana, budaya]	0,176	..	Positif	Negatif	Positif
2	[indonesia, malang, taman, museum, ...]	0,176	..	Positif	Negatif	Negatif

No	Teks Ulasan	indonesia	..	Atraksi	Aksesibilitas	Amenitas
3	[alam, marah, lingkungan, tanggung, ...]	0	..	Negatif	Positif	Positif

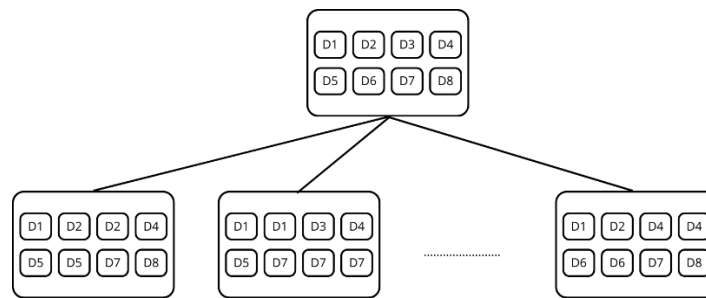
3.6 Algoritma *Random Forest*

Pada penelitian ini, Random Forest digunakan sebagai algoritma klasifikasi untuk menentukan polaritas sentimen berdasarkan fitur TF-IDF. Implementasi model dilakukan melalui empat tahapan utama yaitu :

1. *Bootstrap Sampling*
2. *Random feature selection*
3. Pembentukan pohon keputusan
4. Mekanisme klasifikasi menggunakan *majority voting*

3.6.1 *Bootstrap Sampling*

Bootstrap sampling merupakan tahapan awal dalam algoritma *Random Forest* yang bertujuan untuk menciptakan variasi pada data pelatihan. Dari dataset pelatihan D dengan jumlah data sebanyak N , dilakukan proses pengambilan sampel secara acak dengan pengembalian (sampling with replacement). Proses ini menghasilkan T_{subset} data yang masing-masing memiliki ukuran yang sama dengan dataset awal, namun dengan komposisi data yang berbeda. Implementasi visual dari teknik bootstrap sampling dapat dilihat pada Gambar 3.6.



Gambar 3.6 Simulasi *Bootstrap Sampling*

Pada Gambar 3.6, dataset awal terdiri dari delapan data (D1 hingga D8) yang ditunjukkan pada kotak bagian atas. Dari dataset tersebut, dilakukan pengambilan sampel secara acak dengan pengembalian sebanyak n kali ($n = 8$) untuk membentuk setiap subset bootstrap. Pada subset pertama, data D2, D5, dan D7 terpilih lebih dari satu kali, sedangkan beberapa data lain seperti D6 tidak terpilih sama sekali. Pada subset kedua, data D1 dan D7 muncul berulang, sementara D2, D6, dan D8 tidak ikut diambil. Pola serupa terjadi pada subset-subset berikutnya, di mana setiap subset memiliki kombinasi data yang berbeda akibat proses pengambilan acak tersebut.

Meskipun komposisinya berbeda-beda, setiap subset tetap memiliki ukuran yang sama dengan dataset awal (delapan data), namun memungkinkan adanya data yang terpilih lebih dari satu kali sementara data lain tidak terpilih sama sekali (disebut out-of-bag data). Setiap subset bootstrap selanjutnya digunakan untuk membangun satu pohon keputusan secara independen. Proses pengulangan ini dilakukan sebanyak jumlah pohon yang diinginkan dalam Random Forest, sehingga menghasilkan kumpulan pohon yang masing-masing dilatih dengan data yang

sedikit berbeda. Variasi data inilah yang berkontribusi dalam mengurangi variansi model dan meningkatkan stabilitas prediksi secara keseluruhan.

3.6.2 *Random Feature Selection*

Dalam implementasi algoritma *Random Forest*, digunakan mekanisme seleksi fitur acak di mana m hanya sejumlah fitur yang dipilih secara acak dari total M fitur hasil pembobotan TF-IDF, dengan syarat $m < M$. Proses seleksi ini dilakukan secara repetitif pada setiap *node* selama tahap pemisahan (*splitting*), bukan hanya di awal pembentukan pohon tunggal. Dalam konteks dataset ulasan wisata ini, jika total *vocabulary* yang dihasilkan dari tahap TF-IDF berjumlah 13 fitur seperti pada Tabel 3.8, maka implementasi dari persamaan 2.5 maka sistem hanya akan memilih sekitar 3 hingga 4 fitur secara acak pada setiap *node* saat membangun pohon keputusan.

3.6.3 *Pembentukan Pohon Keputusan*

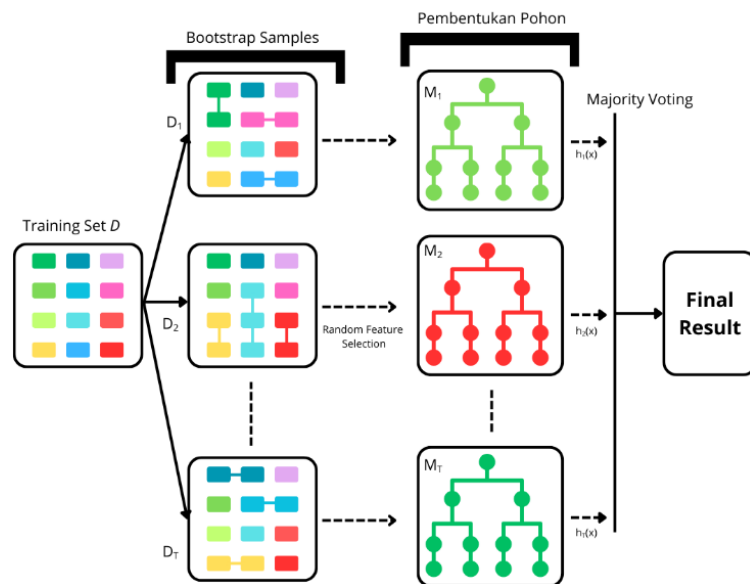
Setiap pohon keputusan dalam model *Random Forest* dibangun menggunakan *subset* data yang diperoleh melalui proses *bootstrap* sampling. Pendekatan ini memungkinkan setiap pohon memiliki variasi distribusi data sehingga meningkatkan keberagaman model. Penggunaan kriteria *Gini Impurity* memungkinkan model memilih atribut yang mampu memisahkan data secara lebih optimal pada setiap *node*. Dengan kombinasi *bootstrap* sampling dan perhitungan *Gini Impurity*, setiap pohon dalam *Random Forest* dapat terbentuk secara beragam namun tetap mempertahankan kemampuan klasifikasi yang akurat. Implementasi ilustrasi visual dari pembentukan pohon dapat dilihat pada Gambar 3.7 berikut :

3.6.4 Mekanisme Klasifikasi (Majority Voting)

Setelah seluruh pohon keputusan dalam model *Random Forest* terbentuk, tahap selanjutnya adalah proses klasifikasi terhadap data uji. Setiap pohon akan menerima data masukan yang sama dan menghasilkan satu prediksi kelas berdasarkan struktur keputusan yang telah dipelajari. Untuk menentukan keluaran akhir, digunakan mekanisme *majority voting* sebagai metode agregasi prediksi.

Dalam mekanisme ini, jika terdapat T pohon dan masing-masing pohon menghasilkan prediksi $h_t(x)$, berdasarkan persamaan 2.7 maka penentuan hasil keputusan akhir didasarkan pada *mode* yaitu suara terbanyak dari prediksi masing-masing pohon keputusan. Pendekatan ini memungkinkan model mengurangi kesalahan individual dari masing-masing pohon keputusan. Dengan demikian, hasil klasifikasi yang diperoleh menjadi lebih konsisten dan memiliki tingkat generalisasi yang lebih baik dibandingkan penggunaan satu pohon tunggal.

Keempat tahapan di atas, yaitu bootstrap sampling, random feature selection, pembentukan pohon keputusan, dan majority voting, bekerja secara terintegrasi membentuk model Random Forest yang robust. Ilustrasi keseluruhan alur tersebut dapat dilihat pada Gambar 3.7 berikut.



Gambar 3.7 Ilustrasi Pembentukan Pohon Keputusan

Berdasarkan Gambar 3.7 mengilustrasikan keseluruhan alur kerja algoritma Random Forest secara terpadu. Training Set D digunakan sebagai sumber data awal yang kemudian menghasilkan beberapa subset D_1, D_2 , hingga D_T melalui proses bootstrap sampling. Setiap subset selanjutnya digunakan untuk membangun satu pohon keputusan secara independen dengan menerapkan random feature selection pada setiap node pemisahan, menghasilkan pohon M_1, M_2 , hingga M_T . Masing-masing pohon kemudian menghasilkan prediksi $h_1(x), h_2(x)$, hingga $h_T(x)$ yang selanjutnya diagregasi melalui mekanisme majority voting untuk menghasilkan Final Result sebagai keluaran klasifikasi akhir.

3.7 Evaluasi Pengujian

Tahapan berikutnya adalah evaluasi pengujian guna mengukur efektivitas kinerja sistem. Instrumen yang digunakan dalam penilaian ini adalah *confusion matrix*, yang mengevaluasi performa melalui empat parameter utama: akurasi, presisi, *recall*, dan *f1-score*. Secara khusus, nilai akurasi berfungsi sebagai indikator untuk menentukan sejauh mana sistem mampu mengklasifikasikan ulasan secara tepat (Amalia Permadi, 2020). Dalam penelitian ini, proses pengujian secara hierarkis melalui dua tingkatan klasifikasi yang bersifat independen satu sama lain. Pada klasifikasi tingkat pertama, sistem melakukan identifikasi keberadaan aspek pada setiap ulasan dengan skema klasifikasi biner "Sentimen" dan "Non-Sentimen", di mana label "Netral" dan "None" digabungkan ke dalam kategori "Non-Sentimen" untuk memfokuskan model pada ulasan yang memiliki relevansi aspek. Pada klasifikasi tingkat kedua, model dilatih dan diuji secara independen menggunakan subset data yang hanya memuat ulasan berlabel "Positif" atau "Negatif" berdasarkan *ground truth*. Pendekatan ini memastikan bahwa performa klasifikasi sentimen dapat diukur secara murni tanpa dipengaruhi oleh kesalahan prediksi pada tahap klasifikasi aspek sebelumnya.

Evaluasi klasifikasi aspek bertujuan untuk mengukur kemampuan model dalam mendeteksi keberadaan suatu aspek pada setiap ulasan. Seluruh data sejumlah 4.195 ulasan digunakan pada tahap ini, di mana setiap ulasan diberikan label biner secara independen untuk masing-masing aspek. Label "Positif" dan "Negatif" dikonversi menjadi kelas "Sentimen" (1), sedangkan label "Netral" dan "None" dikonversi menjadi kelas "Non-Sentimen" (0). Proses konversi ini

dilakukan secara programatik sebelum pelatihan model sehingga setiap aspek memiliki distribusi label binernya masing-masing. Struktur confusion matrix yang digunakan pada tahap ini disajikan pada Tabel 3.14 berikut.

Tabel 3.14 *Confusion Matrix Aspect*

Kelas Aktual	Kelas Prediksi	
	Sentimen	Non Sentimen
Sentimen	20	10
Non Sentimen	10	30

Untuk mengevaluasi *confusion matrix* berdasarkan aspek yang memiliki 2 kelas, maka perhitungan untuk mencari nilai TP-TN-FP-FN dilakukan pada tiap kelasnya. Berikut perhitungan nilai parameter *confusion matrix* pada data di atas:

1. Akurasi

Akurasi merupakan rasio prediksi benar dibandingkan dengan total keseluruhan data.

$$Akurasi = \frac{TP + TN}{Jumlah\ Data} = \frac{50}{70} = 0,71$$

2. Presisi

Presisi dihitung untuk setiap aspek, kemudian dirata-ratakan menggunakan

Macro-average :

$$Presisi_{Atraksi} = \frac{TP}{TP + FP} = \frac{20}{30} = 0,66$$

3. Recall

Recall dihitung untuk setiap aspek, kemudian dirata-ratakan menggunakan

Macro-average :

$$Recall_{Atraksi} = \frac{TP}{TP + FN} = \frac{20}{30} = 0,66$$

4. F1- Score

$$F1 - Score = 2 \times \frac{Presisi \times Recall}{Presisi + Recall} = 2 \times \frac{0,66 \times 0,66}{0,66 + 0,66} = 0,66$$

Evaluasi klasifikasi sentimen bertujuan untuk mengukur kemampuan model dalam menentukan polaritas sentimen suatu ulasan. Berbeda dengan tahap sebelumnya, model pada tahap ini dilatih dan diuji menggunakan subset data yang hanya memuat ulasan berlabel "Positif" atau "Negatif" berdasarkan *ground truth* masing-masing aspek. Pendekatan ini menyebabkan jumlah data yang digunakan pada tahap klasifikasi sentimen berbeda untuk setiap aspek, bergantung pada distribusi label aktual dalam dataset. Evaluasi tetap dilakukan secara independen untuk masing-masing aspek dengan skema klasifikasi biner "Positif" dan "Negatif". Struktur *confusion matrix* yang digunakan pada tahap ini disajikan pada Tabel 3.15 berikut.

Tabel 3.15 *Confusion Matrix Sentiment*

Kelas Aktual	Kelas Prediksi	
	Positif	Negatif
Positif	20	3
Negatif	4	15

Berdasarkan data pada Tabel 3.15, didapatkan nilai parameter *confusion matrix* sebagai berikut:

- a. TP (*True Positive*) = 20
- b. TN (*True Negative*) = 15
- c. FP (*False Positive*) = 4
- d. FN (*False Negative*) = 3

Berikut adalah hasil perhitungan empat metrik evaluasi untuk klasifikasi sentimen:

1. Akurasi (*Accuracy*)

Berdasarkan persamaan 2.8, perhitungan akurasi adalah:

$$Akurasi = \frac{20 + 15}{20 + 15 + 4 + 3} = \frac{35}{42} = 0,833$$

2. Presisi (*Precision*)

Berdasarkan persamaan 2.9 , perhitungan presisi adalah:

$$Presisi = \frac{20}{20 + 4} = \frac{20}{24} = 0,833$$

3. *Recall*

Berdasarkan Persamaan 2.10, perhitungan *recall* adalah:

$$Recall = \frac{20}{20 + 3} = \frac{20}{23} = 0,869$$

4. *F1-Score*

Berdasarkan persamaan 2.11, perhitungan *F1-Score* adalah:

$$F1 - Score = 2 \times \frac{0,833 \times 0,869}{0,833 + 0,869} = \frac{1,448}{1,702} = 0,851$$

3.8 Skenario Pengujian

Tahap berikutnya adalah penyusunan skenario pengujian guna mengevaluasi performa sistem dalam melakukan klasifikasi. Proses ini diawali dengan membagi dataset ke dalam dua bagian: data latih (training) untuk membangun model dan data uji (testing) untuk memvalidasi model tersebut. Pembagian data dilakukan secara

terpisah untuk setiap tahap klasifikasi sesuai dengan subset data yang digunakan pada masing-masing tahap.

Selanjutnya, pengujian dilakukan dengan membandingkan beberapa variasi skenario yang mencakup perbedaan rasio pembagian data dan jumlah pohon ($n_estimators$), untuk mengidentifikasi kombinasi parameter yang menghasilkan performa klasifikasi paling optimal dan stabil. Detail mengenai variasi skenario yang digunakan dalam penelitian ini dipaparkan pada Tabel 3.16.

Tabel 3.16 Rasio Pembagian Data

Skenario	Presentase Data Latih	Presentase Data Uji	Jumlah Pohon ($n_estimators$)
1	70%	30%	100
2	70%	30%	200
3	80%	20%	100
4	80%	20%	200

Tabel 3.16 menyajikan rincian variasi skenario pengujian yang mencakup rasio pembagian *dataset* dan jumlah pohon yang akan diterapkan dalam penelitian ini. Setelah *dataset* dibagi sesuai skenario yang telah ditentukan, proses dilanjutkan dengan tahap evaluasi untuk membandingkan performa antar skenario. Model klasifikasi dibangun menggunakan algoritma Random Forest dengan memvariasikan parameter $n_estimators$, sedangkan evaluasi performa dilakukan menggunakan *confusion matrix* untuk menghasilkan nilai akurasi, presisi, recall, dan F1-score guna mengukur stabilitas dan konsistensi kinerja model pada setiap kombinasi parameter yang diuji.

Proses pengujian klasifikasi dalam penelitian ini dilakukan melalui dua tahapan yang berjalan secara independen untuk menghasilkan analisis sentimen yang spesifik pada setiap aspek ulasan. Tahap pertama adalah klasifikasi aspek,

yang bertujuan mendeteksi keberadaan muatan opini pada setiap ulasan berdasarkan komponen 3A dengan keluaran kelas Sentimen dan Non-Sentimen. Tahap kedua adalah klasifikasi sentimen, yang berfungsi menentukan polaritas opini menjadi Positif atau Negatif pada ulasan yang secara eksplisit berlabel Positif atau Negatif berdasarkan data anotasi. Kedua tahap dilatih dan dievaluasi secara terpisah menggunakan *subset* data masing-masing. Berikut Tabel 3.17 menunjukkan proses klasifikasi yang digunakan dalam penelitian ini :

Tabel 3.17 Tahapan Klasifikasi Dua Tingkat

Tahap	Nama Proses	Deskripsi Tugas	Output Kelas
1	Klasifikasi Aspek	Mengidentifikasi kategori topik dalam ulasan berdasarkan komponen 3A.	Sentimen, Non Sentimen
2	Klasifikasi Sentimen	Menentukan polaritas opini untuk masing-masing aspek yang terdeteksi.	Positif, Negatif.

Berdasarkan rincian pada Tabel 3.17, tahap pertama merupakan klasifikasi aspek yang bertujuan mendeteksi keberadaan muatan opini pada setiap ulasan untuk masing-masing komponen 3A pariwisata, yaitu Atraksi, Aksesibilitas, dan Amenitas, dengan keluaran biner Sentimen dan Non-Sentimen. Tahap kedua merupakan klasifikasi sentimen yang bertujuan menentukan polaritas opini menjadi Positif atau Negatif. Kedua tahap dalam penelitian ini dilatih dan dievaluasi secara independen menggunakan *subset* masing-masing, tahap klasifikasi aspek menggunakan seluruh *dataset* sedangkan tahap klasifikasi sentimen menggunakan *subset* data yang berlabel Positif atau Negatif. Pendekatan ini dipilih untuk mengukur kemampuan setiap tahap secara terpisah sehingga performa yang

dilaporkan mencerminkan potensi model pada masing-masing komponen klasifikasi.

BAB IV

HASIL DAN PEMBAHASAN

4.1 Hasil Uji Coba

Tahapan ini memaparkan hasil evaluasi klasifikasi yang dilakukan oleh sistem melalui perbandingan antara data aktual dan data prediksi menggunakan instrumen *confusion matrix*. Melalui perbandingan tersebut diperoleh nilai performa berupa akurasi, presisi, *recall*, dan *F1-score* dari model yang telah dikembangkan. Pengukuran metrik evaluasi difokuskan pada kelas target penelitian, yaitu kelas Sentimen pada tahap Klasifikasi Aspek dan kelas Positif pada tahap Klasifikasi Sentimen, sesuai dengan tujuan penelitian yang berfokus pada klasifikasi sentimen ulasan wisata.

Pengujian dilakukan melalui dua tahap klasifikasi yang berjalan secara independen meliputi aspek Atraksi, Aksesibilitas, dan Amenitas. Pada tahap pertama yaitu Klasifikasi Aspek, model mendeteksi keberadaan muatan opini pada setiap ulasan dengan skema biner "Sentimen" dan "Non-Sentimen", di mana label "Netral" dan "None" digabungkan ke dalam kategori "Non-Sentimen" karena keduanya tidak mengandung sinyal polaritas yang dapat diklasifikasikan. Pada tahap kedua yaitu Klasifikasi Sentimen, model menentukan polaritas opini menjadi "Positif" atau "Negatif" menggunakan *subset* data yang berlabel Positif atau Negatif. Kedua tahap dilatih dan dievaluasi secara terpisah menggunakan *subset* data masing-masing.

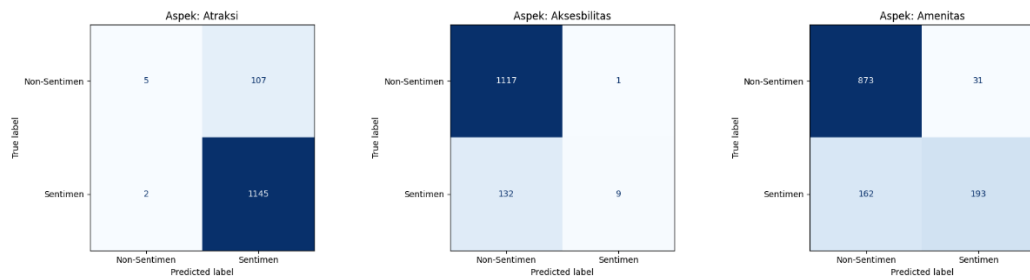
Seluruh proses pengujian dilakukan pada empat skenario yang memvariasikan rasio pembagian data latih dan data uji sebesar 70:30 serta 80:20, dikombinasikan dengan variasi parameter $n_estimators$ sebesar 100 dan 200, untuk mendapatkan hasil analisis performa model klasifikasi pada setiap kombinasi parameter yang diuji.

4.1.1 Hasil Uji Skenario 1

Pada skenario pertama, pengujian dilakukan menggunakan proporsi pembagian data sebesar 70% untuk data latih dan 30% untuk data uji dengan parameter $n_estimators$ sebesar 100. *Dataset* yang digunakan terdiri dari 4.195 ulasan yang menghasilkan 7.432 fitur kata setelah melalui proses ekstraksi TF-IDF. Pada tahap Klasifikasi Aspek, seluruh *dataset* digunakan sehingga sebanyak 2.936 ulasan digunakan sebagai data latih dan 1.259 ulasan digunakan sebagai data uji.

Pada tahap Klasifikasi Sentimen, data disiapkan secara independen dengan memfilter *subset* yang berlabel Positif atau Negatif dari *dataset* awal untuk masing-masing aspek, kemudian dibagi dengan rasio yang sama dengan pengujian dilakukan melalui dua tahap klasifikasi secara independen.

Tahap pertama adalah Klasifikasi Aspek, pada tahap ini model melakukan klasifikasi biner untuk mendeteksi keberadaan muatan opini pada setiap ulasan. Kelas "Sentimen" merepresentasikan ulasan yang mengandung muatan opini pada suatu aspek, sedangkan kelas "Non-Sentimen" merepresentasikan ulasan yang tidak mengandung opini relevan pada aspek tersebut. Hasil klasifikasi untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.1 berikut.

Gambar 4.1 *Confusion Matrix* Klasifikasi Aspek Skenario 1

Berdasarkan *confusion matrix* pada Gambar 4.1, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.1 berikut :

Tabel 4.1 *Confusion Matrix* Klasifikasi Aspek Skenario 1

Aspek	TP	TN	FP	FN	Total
Atraksi	1.145	5	107	2	1259
Aksesibilitas	9	1117	1	132	1259
Amenitas	193	873	31	162	1259

Berdasarkan Tabel 4.1, dilakukan perhitungan metrik evaluasi secara manual. Sebagai representasi, berikut disajikan perhitungan untuk aspek Amenitas yang memiliki distribusi data paling berimbang di antara ketiga aspek.

$$\text{Akurasi}_{\text{Amenitas}} = \frac{TP + TN}{N} = \frac{193 + 873}{1259} = \frac{1066}{1259} = 0,85$$

$$\text{Presisi}_{\text{Amenitas}} = \frac{TP}{TP + FP} = \frac{193}{193 + 31} = \frac{193}{224} = 0,86$$

$$\text{Recall}_{\text{Amenitas}} = \frac{TP}{TP + FN} = \frac{193}{193 + 162} = \frac{193}{355} = 0,54$$

$$F1 - \text{Score}_{\text{Amenitas}} = \frac{2 \times \text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}} = \frac{2 \times 0,86 \times 0,54}{0,86 + 0,54} = 0,67$$

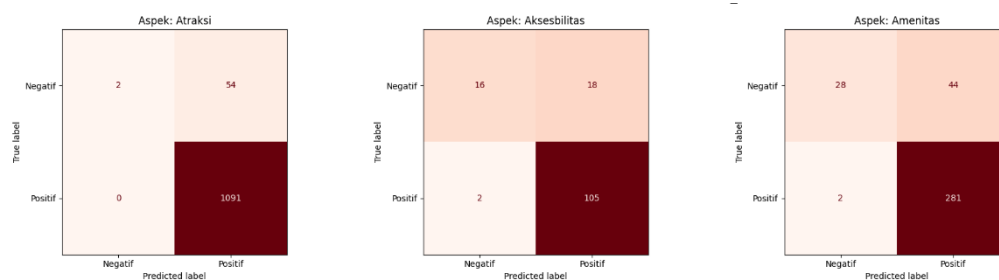
Hasil perhitungan tersebut konsisten dengan nilai yang tercantum pada *output* evaluasi model. Nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Aspek secara keseluruhan disajikan pada Tabel 4.2 berikut.

Tabel 4.2 Evaluasi Metrik Klasifikasi Aspek Skenario 1

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,91	0,91	0,99	0,95
Aksesibilitas	0,89	0,90	0,06	0,12
Amenitas	0,85	0,86	0,54	0,67
Rata-rata	0,88	0,89	0,53	0,58

Berdasarkan Tabel 4.2, aspek Atraksi memperoleh F1-score tertinggi sebesar 0,95, diikuti Amenitas sebesar 0,67, sementara Aksesibilitas memperoleh nilai terendah sebesar 0,12. Secara rata-rata sistem Klasifikasi Aspek pada skenario 1 memperoleh akurasi 0,88, presisi 0,89, *recall* 0,53, dan F1-score 0,58.

Tahap kedua adalah klasifikasi sentimen, pada tahap ini model mengklasifikasikan polaritas opini pada ulasan menjadi "Positif" atau "Negatif". Data yang digunakan disiapkan secara independen dengan memfilter langsung dari *dataset* awal hanya ulasan yang berlabel Positif atau Negatif berdasarkan hasil anotasi, kemudian dibagi dengan rasio yang sama yaitu 70% data latih dan 30% data uji. Jumlah data yang tersedia berbeda untuk setiap aspek, yaitu aspek Atraksi sebanyak 3.822 data, Aksesibilitas sebanyak 470 data, dan Amenitas sebanyak 1.183 data. Hasil klasifikasi untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.2 berikut.

Gambar 4.2 *Confusion Matrix* Klasifikasi Sentimen Skenario 1

Berdasarkan *confusion matrix* pada Gambar 4.2, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.3 berikut

Tabel 4.3 Confusion Matrix Klasifikasi Sentimen Skenario 1

Aspek	TP	TN	FP	FN	Total
Atraksi	1091	2	54	0	1147
Aksesibilitas	105	16	18	2	141
Amenitas	281	28	44	2	355

Berdasarkan Tabel 4.3, dilakukan perhitungan metrik evaluasi secara manual menggunakan aspek Aksesibilitas sebagai representasi karena memiliki distribusi kelas paling berimbang pada tahap Klasifikasi Sentimen.

$$\text{Akurasi}_{\text{Aksesibilitas}} = \frac{TP + TN}{N} = \frac{105 + 16}{141} = \frac{121}{141} = 0,86$$

$$\text{Presisi}_{\text{Aksesibilitas}} = \frac{TP}{TP + FP} = \frac{105}{105 + 18} = \frac{105}{123} = 0,85$$

$$\text{Recall}_{\text{Aksesibilitas}} = \frac{TP}{TP + FN} = \frac{105}{105 + 2} = \frac{105}{107} = 0,98$$

$$F1 - \text{Score}_{\text{Aksesibilitas}} = \frac{2 \times \text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}} = \frac{2 \times 0,85 \times 0,98}{0,85 + 0,98} = 0,67$$

Hasil perhitungan tersebut konsisten dengan nilai yang tercantum pada *output* evaluasi model. Nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Sentimen secara keseluruhan disajikan pada Tabel 4.4 berikut.

Tabel 4.4 Evaluasi Metrik Klasifikasi Sentimen Skenario 1

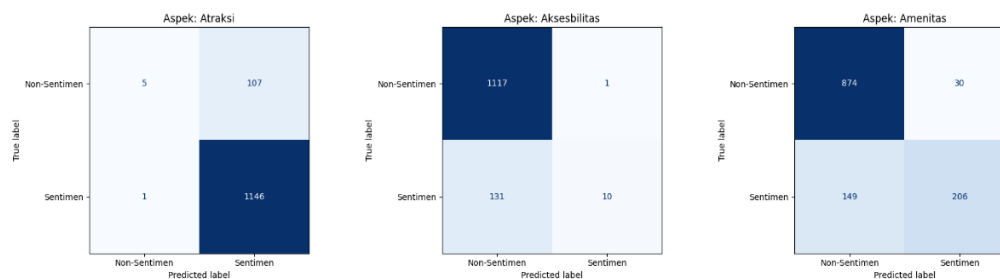
Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,95	0,95	1,00	0,97
Aksesibilitas	0,85	0,85	0,98	0,91
Amenitas	0,87	0,86	0,99	0,92
Rata-rata	0,89	0,89	0,99	0,93

Berdasarkan Tabel 4.4, sistem Klasifikasi Sentimen pada skenario 1 secara keseluruhan memperoleh rata-rata akurasi sebesar 0,89, presisi sebesar 0,89, *recall* sebesar 0,99, dan *F1-score* sebesar 0,93. Aspek Atraksi memperoleh F1-score tertinggi sebesar 0,97, diikuti Amenitas sebesar 0,92, sementara Aksesibilitas memperoleh nilai sebesar 0,91.

4.1.2 Hasil Uji Skenario 2

Pada skenario kedua, pengujian dilakukan menggunakan proporsi pembagian data sebesar 70% untuk data latih dan 30% untuk data uji dengan parameter *n_estimators* sebesar 200. *Dataset* yang digunakan terdiri dari 4.195 ulasan yang menghasilkan 7.432 fitur kata setelah melalui proses ekstraksi TF-IDF. Pada tahap Klasifikasi Aspek, seluruh *dataset* digunakan sehingga sebanyak 2.936 ulasan digunakan sebagai data latih dan 1.259 ulasan digunakan sebagai data uji. Pada tahap Klasifikasi Sentimen, data disiapkan secara independen dengan memfilter *subset* yang berlabel Positif atau Negatif dari *dataset* awal untuk masing-masing aspek, kemudian dibagi dengan rasio yang sama dengan pengujian dilakukan melalui dua tahap klasifikasi secara independen.

Tahap pertama adalah Klasifikasi Aspek, pada tahap ini model melakukan klasifikasi biner untuk mendeteksi keberadaan muatan opini pada setiap ulasan. Kelas "Sentimen" merepresentasikan ulasan yang mengandung muatan opini pada suatu aspek, sedangkan kelas "Non-Sentimen" merepresentasikan ulasan yang tidak mengandung opini relevan pada aspek tersebut. Hasil klasifikasi untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.3 berikut.

Gambar 4.3 *Confusion Matrix* Klasifikasi Aspek Skenario 2

Berdasarkan *confusion matrix* pada Gambar 4.3, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.5 berikut :

Tabel 4.5 *Confusion Matrix* Klasifikasi Aspek Skenario 2

Aspek	TP	TN	FP	FN	Total
Atraksi	1.146	5	107	1	1259
Aksesibilitas	10	1117	1	131	1259
Amenitas	206	874	30	149	1259

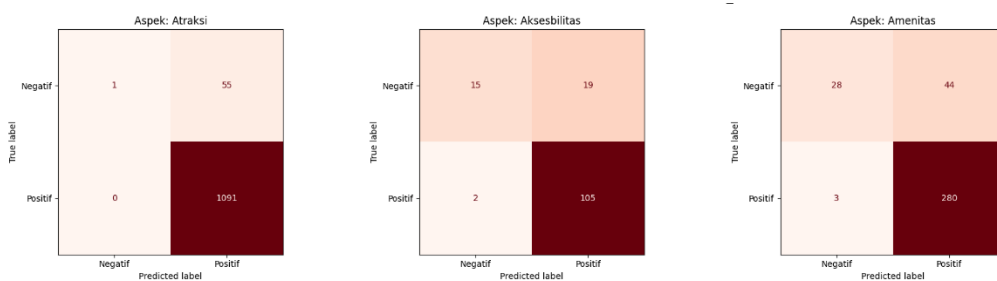
Berdasarkan Tabel 4.5, nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Aspek secara keseluruhan disajikan pada Tabel 4.6 berikut :

Tabel 4.6 Evaluasi Metrik Klasifikasi Aspek Skenario 2

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,91	0,91	0,99	0,95
Aksesibilitas	0,90	0,91	0,07	0,13
Amenitas	0,86	0,87	0,58	0,70
Rata-rata	0,89	0,90	0,55	0,59

Berdasarkan Tabel 4.6, aspek Atraksi memperoleh F1-score tertinggi sebesar 0,95, diikuti Amenitas sebesar 0,70, sementara Aksesibilitas memperoleh nilai terendah sebesar 0,13. Secara rata-rata sistem Klasifikasi Aspek pada skenario 2 memperoleh akurasi 0,89, presisi 0,90, *recall* 0,55, dan *F1-score* 0,59.

Tahap kedua adalah klasifikasi sentimen, pada tahap ini model mengklasifikasikan polaritas opini pada ulasan menjadi "Positif" atau "Negatif". Data yang digunakan disiapkan secara independen dengan memfilter langsung dari *dataset* awal hanya ulasan yang berlabel Positif atau Negatif berdasarkan hasil anotasi, kemudian dibagi dengan rasio yang sama yaitu 70% data latih dan 30% data uji. Jumlah data yang tersedia berbeda untuk setiap aspek, yaitu aspek Atraksi sebanyak 3.822 data, Aksesibilitas sebanyak 470 data, dan Amenitas sebanyak 1.183 data. Hasil klasifikasi untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.4 berikut.



Gambar 4.4 *Confusion Matrix* Klasifikasi Sentimen Skenario 2

Berdasarkan *confusion matrix* pada Gambar 4.4, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.7 berikut.

Tabel 4.7 *Confusion Matrix* Klasifikasi Sentimen Skenario 2

Aspek	TP	TN	FP	FN	Total
Atraksi	1091	1	55	0	1147
Aksesibilitas	105	15	19	2	141
Amenitas	280	28	44	3	355

Berdasarkan Tabel 4.7, nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Sentimen secara keseluruhan disajikan pada Tabel 4.8 berikut.

Tabel 4.8 Evaluasi Metrik Klasifikasi Sentimen Skenario 2

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,95	0,95	1,00	0,97
Aksesibilitas	0,85	0,85	0,98	0,91
Amenitas	0,87	0,86	0,99	0,92
Rata-rata	0,89	0,89	0,99	0,93

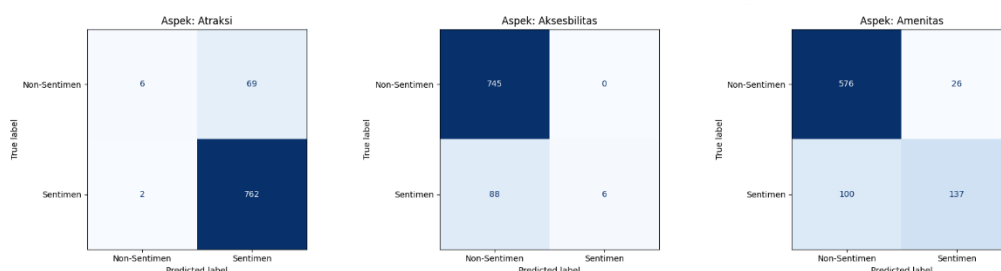
Berdasarkan Tabel 4.8, aspek Atraksi memperoleh F1-score tertinggi sebesar 0,97, diikuti Amenitas sebesar 0,92, sementara Aksesibilitas memperoleh nilai terendah sebesar 0,91. Secara rata-rata sistem Klasifikasi Sentimen pada skenario 2 memperoleh akurasi 0,89, presisi 0,89, *recall* 0,99, dan *F1-score* 0,93.

4.1.3 Hasil Uji Skenario 3

Pada skenario ketiga, pengujian dilakukan menggunakan proporsi pembagian data sebesar 80% untuk data latih dan 20% untuk data uji dengan parameter *n_estimators* sebesar 100. *Dataset* yang digunakan terdiri dari 4.195 ulasan yang menghasilkan 7.432 fitur kata setelah melalui proses ekstraksi TF-IDF. Pada tahap Klasifikasi Aspek, seluruh *dataset* digunakan sehingga sebanyak 3.356 ulasan digunakan sebagai data latih dan 839 ulasan digunakan sebagai data uji.

Pada tahap Klasifikasi Sentimen, data disiapkan secara independen dengan memfilter *subset* yang berlabel Positif atau Negatif dari *dataset* awal untuk masing-masing aspek, kemudian dibagi dengan rasio yang sama dengan pengujian dilakukan melalui dua tahap klasifikasi secara independen.

Hasil klasifikasi tahap pertama Klasifikasi Aspek untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.5 berikut.



Gambar 4.5 *Confusion Matrix* Klasifikasi Aspek Skenario 3

Berdasarkan *confusion matrix* pada Gambar 4.5, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.9 berikut :

Tabel 4.9 *Confusion Matrix* Klasifikasi Aspek Skenario 3

Aspek	TP	TN	FP	FN	Total
Atraksi	762	6	69	2	839
Aksesibilitas	6	745	0	88	839
Amenitas	137	576	26	100	839

Berdasarkan Tabel 4.9, nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Aspek secara keseluruhan disajikan pada Tabel 4.10 berikut :

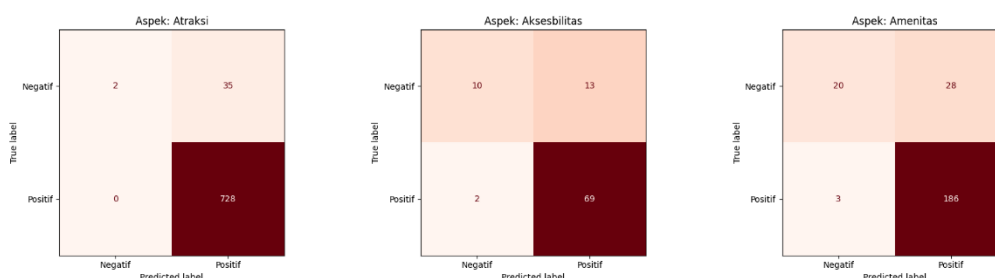
Tabel 4.10 Evaluasi Metrik Klasifikasi Aspek Skenario 3

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,92	0,92	0,99	0,95
Aksesibilitas	0,90	1,00	0,06	0,12
Amenitas	0,85	0,84	0,58	0,69
Rata-rata	0,89	0,92	0,54	0,59

Berdasarkan Tabel 4.10, aspek Atraksi memperoleh F1-score tertinggi sebesar 0,95, diikuti Amenitas sebesar 0,69, sementara Aksesibilitas memperoleh

nilai terendah sebesar 0,12. Secara rata-rata sistem Klasifikasi Aspek pada skenario 3 memperoleh akurasi 0,89, presisi 0,92, *recall* 0,54, dan *F1-score* 0,59.

Tahap kedua klasifikasi sentimen dengan jumlah data yang tersedia berbeda untuk setiap aspek, yaitu aspek Atraksi sebanyak 3.822 data, Aksesibilitas sebanyak 470 data, dan Amenitas sebanyak 1.183 data. Hasil klasifikasi tahap kedua Klasifikasi Sentimen untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.6 berikut.



Gambar 4.6 *Confusion Matrix* Klasifikasi Sentimen Skenario 3

Berdasarkan *confusion matrix* pada Gambar 4.6, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.11 berikut.

Tabel 4.11 *Confusion Matrix* Klasifikasi Sentimen Skenario 3

Aspek	TP	TN	FP	FN	Total
Atraksi	728	2	35	0	765
Aksesibilitas	69	10	13	2	94
Amenitas	186	20	28	3	237

Berdasarkan Tabel 4.11, nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Sentimen secara keseluruhan disajikan pada Tabel 4.12 berikut.

Tabel 4.12 Evaluasi Metrik Klasifikasi Sentimen Skenario 3

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,95	0,95	1,00	0,97

Aspek	Akurasi	Presisi	Recall	F1-Score
Aksesibilitas	0,84	0,84	0,97	0,90
Amenitas	0,87	0,87	0,98	0,92
Rata-rata	0,89	0,89	0,98	0,93

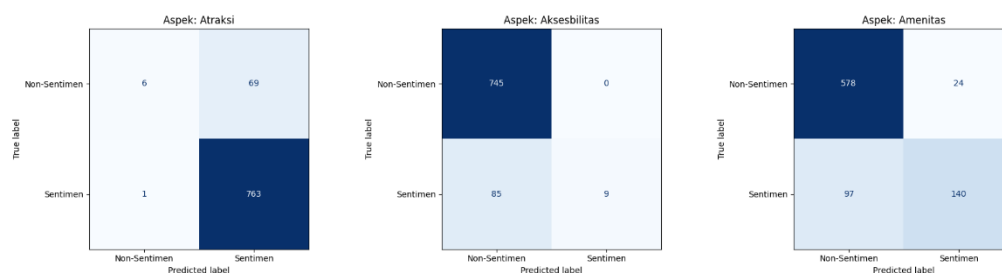
Berdasarkan Tabel 4.12, aspek Atraksi memperoleh *F1-score* tertinggi sebesar 0,97, diikuti Amenitas sebesar 0,92, sementara Aksesibilitas memperoleh nilai terendah sebesar 0,90. Secara rata-rata sistem Klasifikasi Sentimen pada skenario 3 memperoleh akurasi 0,89, presisi 0,89, *recall* 0,98, dan *F1-score* 0,93.

4.1.4 Hasil Uji Skenario 4

Pada skenario keempat, pengujian dilakukan menggunakan proporsi pembagian data sebesar 80% untuk data latih dan 20% untuk data uji dengan parameter *n_estimators* sebesar 200. *Dataset* yang digunakan terdiri dari 4.195 ulasan yang menghasilkan 7.432 fitur kata setelah melalui proses ekstraksi TF-IDF. Pada tahap Klasifikasi Aspek, seluruh *dataset* digunakan sehingga sebanyak 3.356 ulasan digunakan sebagai data latih dan 839 ulasan digunakan sebagai data uji.

Pada tahap Klasifikasi Sentimen, data disiapkan secara independen dengan memfilter *subset* yang berlabel Positif atau Negatif dari *dataset* awal untuk masing-masing aspek, kemudian dibagi dengan rasio yang sama dengan pengujian dilakukan melalui dua tahap klasifikasi secara independen.

Hasil klasifikasi tahap pertama Klasifikasi Aspek untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.7 berikut.



Gambar 4.7 *Confusion Matrix* Klasifikasi Aspek Skenario 4

Berdasarkan *confusion matrix* pada Gambar 4.7, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.13 berikut :

Tabel 4.13 *Confusion Matrix* Klasifikasi Aspek Skenario 4

Aspek	TP	TN	FP	FN	Total
Atraksi	763	6	69	1	839
Aksesibilitas	9	745	0	85	839
Amenitas	140	578	24	97	839

Berdasarkan Tabel 4.13, nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Aspek secara keseluruhan disajikan pada Tabel 4.14 berikut :

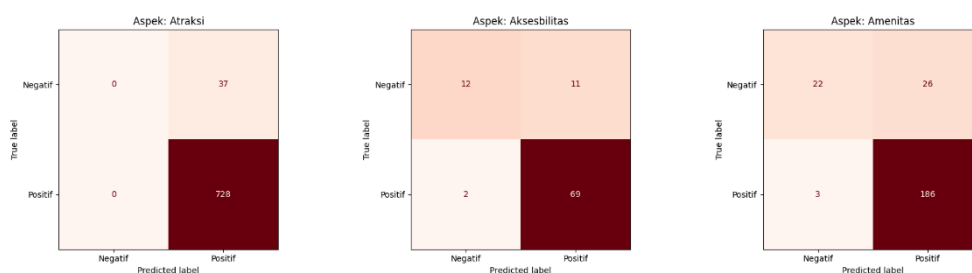
Tabel 4.14 Evaluasi Metrik Klasifikasi Aspek Skenario 4

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,92	0,92	1,00	0,96
Aksesibilitas	0,90	1,00	0,10	0,18
Amenitas	0,86	0,85	0,59	0,70
Rata-rata	0,89	0,92	0,56	0,61

Berdasarkan Tabel 4.14, aspek Atraksi memperoleh F1-score tertinggi sebesar 0,96, diikuti Amenitas sebesar 0,70, sementara Aksesibilitas memperoleh

nilai terendah sebesar 0,18. Secara rata-rata sistem Klasifikasi Aspek pada skenario 4 memperoleh akurasi 0,89, presisi 0,92, *recall* 0,56, dan *F1-score* 0,61.

Tahap kedua klasifikasi sentimen dengan jumlah data yang tersedia berbeda untuk setiap aspek, yaitu aspek Atraksi sebanyak 3.822 data, Aksesibilitas sebanyak 470 data, dan Amenitas sebanyak 1.183 data. Hasil klasifikasi tahap kedua Klasifikasi Sentimen untuk ketiga aspek divisualisasikan melalui *confusion matrix* pada Gambar 4.8 berikut.



Gambar 4.8 *Confusion Matrix* Klasifikasi Sentimen Skenario 4

Berdasarkan *confusion matrix* pada Gambar 4.8, nilai TP, TN, FP, dan FN untuk masing-masing aspek dirangkum pada Tabel 4.15 berikut.

Tabel 4.15 *Confusion Matrix* Klasifikasi Sentimen Skenario 4

Aspek	TP	TN	FP	FN	Total
Atraksi	728	0	37	0	765
Aksesibilitas	69	12	11	2	94
Amenitas	186	22	26	3	237

Berdasarkan Tabel 4.15, nilai metrik untuk seluruh aspek beserta rata-rata sistem Klasifikasi Sentimen secara keseluruhan disajikan pada Tabel 4.16 berikut.

Tabel 4.16 Evaluasi Metrik Klasifikasi Sentimen Skenario 4

Aspek	Akurasi	Presisi	Recall	F1-Score
Atraksi	0,95	0,95	1,00	0,97

Aspek	Akurasi	Presisi	Recall	F1-Score
Aksesibilitas	0,86	0,86	0,97	0,91
Amenitas	0,88	0,88	0,98	0,93
Rata-rata	0,90	0,90	0,98	0,94

Berdasarkan Tabel 4.16, aspek Atraksi memperoleh F1-score tertinggi sebesar 0,97, diikuti Amenitas sebesar 0,93, sementara Aksesibilitas memperoleh nilai terendah sebesar 0,91. Secara rata-rata sistem Klasifikasi Sentimen pada skenario 4 memperoleh akurasi 0,90, presisi 0,90, *recall* 0,98, dan *F1-score* 0,94.

4.2 Pembahasan

Bagian ini memaparkan pembahasan terhadap hasil pengujian yang telah dilakukan pada keempat skenario. Setiap skenario dianalisis berdasarkan nilai metrik evaluasi yang diperoleh, meliputi akurasi, presisi, *recall*, dan *F1-score* pada kedua tahap klasifikasi, yaitu Klasifikasi Aspek dan Klasifikasi Sentimen, untuk masing-masing komponen 3A pariwisata yang terdiri dari aspek Atraksi, Aksesibilitas, dan Amenitas.

4.2.1 Pembahasan Hasil Uji Coba Skenario 1

Pada skenario 1 dengan proporsi pembagian data 70:30 dan parameter *n_estimators* sebesar 100, hasil pengujian menunjukkan perbedaan performa yang cukup signifikan antara tahap Klasifikasi Aspek dan tahap Klasifikasi Sentimen.

Pada tahap Klasifikasi Aspek, aspek Atraksi memperoleh performa terbaik dengan F1-score sebesar 0,95. Hal ini sejalan dengan distribusi data yang sangat didominasi oleh kelas Sentimen, sebagaimana terlihat pada *confusion matrix* di mana nilai TP mencapai 1.145 dari total 1.259 data uji. Kondisi ini menyebabkan

model mudah mengenali pola kelas mayoritas sehingga menghasilkan *recall* yang tinggi sebesar 0,99. Sebaliknya, aspek Aksesibilitas memperoleh F1-score terendah sebesar 0,12 meskipun akurasi mencapai 0,89. Kesenjangan antara akurasi dan F1-score ini disebabkan oleh ketidakseimbangan distribusi kelas yang ekstrem, di mana kelas Non-Sentimen sangat mendominasi dengan 1.117 TN dibanding hanya 9 TP, sehingga model cenderung memprediksi hampir seluruh data sebagai Non-Sentimen dan menghasilkan *recall* yang sangat rendah sebesar 0,06. Aspek Amenitas berada di antara keduanya dengan F1-score sebesar 0,67, didukung oleh distribusi data yang lebih berimbang dibanding dua aspek lainnya. Secara keseluruhan, tahap Klasifikasi Aspek pada skenario 1 memperoleh rata-rata F1-score sebesar 0,58.

Pada tahap Klasifikasi Sentimen, seluruh aspek menunjukkan performa yang lebih baik dibandingkan tahap sebelumnya. Aspek Atraksi memperoleh *F1-score* tertinggi sebesar 0,97 dengan *recall* sempurna 1,00, yang menunjukkan bahwa model mampu mendeteksi seluruh ulasan bersentimen positif pada aspek tersebut tanpa ada yang terlewat. Aspek Amenitas dan Aksesibilitas masing-masing memperoleh *F1-score* sebesar 0,92 dan 0,91 dengan *recall* tinggi di atas 0,98, yang mengindikasikan bahwa model pada tahap ini konsisten dalam mengklasifikasikan polaritas opini. Secara keseluruhan, tahap Klasifikasi Sentimen pada skenario 1 memperoleh rata-rata F1-score sebesar 0,93, jauh lebih tinggi dibandingkan tahap Klasifikasi Aspek.

Perbedaan performa antara kedua tahap ini mengindikasikan bahwa klasifikasi polaritas opini relatif lebih mudah dilakukan oleh model dibandingkan

deteksi keberadaan opini itu sendiri, terutama pada aspek dengan jumlah data yang terbatas seperti Aksesibilitas.

4.2.2 Pembahasan Hasil Uji Coba Skenario 2

Pada skenario 2 dengan proporsi pembagian data 70:30 dan parameter $n_estimators$ sebesar 200, hasil pengujian menunjukkan pola performa yang serupa dengan skenario 1 namun dengan peningkatan kecil pada tahap Klasifikasi Aspek.

Pada tahap Klasifikasi Aspek, aspek Atraksi kembali memperoleh performa terbaik dengan F1-score sebesar 0,95, sama dengan skenario 1. Hal ini mengindikasikan bahwa penambahan jumlah pohon dari 100 menjadi 200 tidak memberikan pengaruh pada aspek yang distribusi datanya sudah sangat didominasi oleh satu kelas. Aspek Aksesibilitas mengalami peningkatan kecil dari *F1-score* 0,12 pada skenario 1 menjadi 0,13 pada skenario 2, dengan *recall* yang sedikit meningkat dari 0,06 menjadi 0,07. Meskipun terjadi peningkatan, nilai tersebut masih sangat rendah dan mencerminkan bahwa penambahan jumlah pohon tidak mampu mengatasi permasalahan ketidakseimbangan distribusi data yang ekstrem pada aspek ini. Aspek Amenitas menunjukkan peningkatan yang lebih terlihat dibanding Aksesibilitas, dengan *F1-score* naik dari 0,67 pada skenario 1 menjadi 0,70 pada skenario 2, didukung oleh peningkatan *recall* dari 0,54 menjadi 0,58. Peningkatan ini mengindikasikan bahwa penambahan jumlah pohon memberikan manfaat yang lebih nyata pada aspek dengan distribusi data yang lebih berimbang. Secara keseluruhan, tahap Klasifikasi Aspek pada skenario 2 memperoleh rata-rata F1-score sebesar 0,59, meningkat tipis dibanding skenario 1 sebesar 0,58.

Pada tahap Klasifikasi Sentimen, hasil skenario 2 menunjukkan performa yang hampir identik dengan skenario 1. Ketiga aspek memperoleh nilai F1-score yang sama yaitu Atraksi 0,97, Amenitas 0,92, dan Aksesibilitas 0,91, dengan rata-rata *F1-score* sistem sebesar 0,93. Perbedaan hanya terlihat pada level *confusion matrix* di mana terdapat perpindahan satu data prediksi pada aspek Atraksi dan Amenitas, namun tidak cukup signifikan untuk mengubah nilai metrik evaluasi secara keseluruhan. Kondisi ini mengindikasikan bahwa pada tahap Klasifikasi Sentimen, model dengan *n_estimators* 100 dan *n_estimators* 200 menghasilkan keputusan klasifikasi yang hampir sepenuhnya sama pada rasio pembagian data 70:30.

Perbedaan performa antara kedua tahap pada skenario 2 tetap konsisten dengan pola yang ditemukan pada skenario 1 tahap Klasifikasi Aspek masih menjadi *bottleneck* dengan rata-rata F1-score 0,59, sementara tahap Klasifikasi Sentimen tetap menunjukkan performa tinggi dengan rata-rata F1-score 0,93. Penambahan *n_estimators* dari 100 menjadi 200 hanya memberikan dampak marginal pada Klasifikasi Aspek dan praktis tidak berpengaruh pada Klasifikasi Sentimen pada rasio pembagian data 70:30 ini.

4.2.3 Pembahasan Hasil Uji Coba Skenario 3

Pada skenario 3 dengan proporsi pembagian data 80:20 dan parameter *n_estimators* sebesar 100, hasil pengujian menunjukkan pola performa yang konsisten dengan dua skenario sebelumnya, meskipun terdapat perbedaan pada jumlah data uji akibat perubahan rasio pembagian data.

Pada tahap Klasifikasi Aspek, aspek Atraksi kembali memperoleh performa terbaik dengan *F1-score* sebesar 0,95, sama dengan skenario 1 dan 2. Konsistensi nilai ini menunjukkan bahwa baik perubahan rasio pembagian data dari 70:30 menjadi 80:20 maupun variasi *n_estimators* tidak memberikan pengaruh pada aspek yang distribusi datanya sangat didominasi oleh satu kelas model tetap mudah mengenali pola kelas Sentimen yang mendominasi dengan *recall* 0,99. Aspek Aksesibilitas kembali memperoleh *F1-score* terendah sebesar 0,12, sama persis dengan skenario 1. Meskipun rasio data latih ditingkatkan menjadi 80%, penambahan data latih tidak mampu memperbaiki performa pada aspek ini karena akar permasalahannya terletak pada ketidakseimbangan distribusi kelas yang ekstrem, bukan pada kurangnya data latih secara absolut. Hal yang menarik pada skenario 3 adalah nilai presisi Aksesibilitas mencapai 1,00 artinya setiap kali model memprediksi Sentimen pada aspek ini, prediksinya selalu benar namun *recall* yang hanya 0,06 menunjukkan bahwa model sangat jarang membuat prediksi Sentimen sama sekali. Aspek Amenitas memperoleh *F1-score* sebesar 0,69, sedikit lebih rendah dibanding skenario 2 yang mencapai 0,70, meskipun menggunakan rasio data latih yang lebih besar. Penurunan kecil ini kemungkinan disebabkan oleh berkurangnya jumlah data uji dari 1.259 menjadi 839 sehingga estimasi metrik menjadi kurang stabil. Secara keseluruhan tahap Klasifikasi Aspek pada skenario 3 memperoleh rata-rata *F1-score* sebesar 0,59, sama dengan skenario 2.

Pada tahap Klasifikasi Sentimen, jumlah data uji per aspek berkurang seiring perubahan rasio menjadi 80:20 maka Atraksi sebanyak 765 data, Aksesibilitas sebanyak 94 data dan Amenitas sebanyak 237 data. Aspek Atraksi kembali

memperoleh *F1-score* tertinggi sebesar 0,97 dengan *recall* sempurna 1,00, konsisten dengan dua skenario sebelumnya. Aspek Amenitas memperoleh *F1-score* sebesar 0,92, sama dengan skenario 1 dan 2. Aspek Aksesibilitas mengalami penurunan kecil dengan *F1-score* 0,90 dibanding 0,91 pada skenario 1 dan 2, dengan *recall* yang sedikit turun dari 0,98 menjadi 0,97. Penurunan ini sejalan dengan berkurangnya jumlah data uji Aksesibilitas dari 141 menjadi 94 data, sehingga estimasi metrik menjadi lebih sensitif terhadap setiap kesalahan prediksi. Secara keseluruhan tahap Klasifikasi Sentimen pada skenario 3 memperoleh rata-rata *F1-score* sebesar 0,93, identik dengan skenario 1 dan 2.

Pola yang muncul dari skenario 3 memperkuat temuan sebelumnya bahwa tahap Klasifikasi Sentimen menunjukkan performa yang sangat stabil di *angka F1-score* 0,93 meskipun rasio pembagian data berubah, sementara tahap Klasifikasi Aspek tetap menjadi *bottleneck* dengan rata-rata *F1-score* 0,59. Perubahan rasio data latih dari 70% menjadi 80% tidak memberikan dampak yang signifikan terhadap performa keseluruhan sistem pada *n_estimators* 100.

4.2.4 Pembahasan Hasil Uji Coba Skenario 4

Pada skenario 4 dengan proporsi pembagian data 80:20 dan parameter *n_estimators* sebesar 200, hasil pengujian menunjukkan peningkatan performa yang lebih terlihat dibanding skenario sebelumnya, baik pada tahap Klasifikasi Aspek maupun Klasifikasi Sentimen.

Pada tahap Klasifikasi Aspek, aspek Atraksi untuk pertama kali memperoleh *recall* sempurna 1,00 dengan *F1-score* 0,96, meningkat dari 0,95 pada tiga skenario sebelumnya. Peningkatan ini terjadi karena kombinasi rasio data latih yang lebih

besar (80%) dan jumlah pohon yang lebih banyak (200) memungkinkan model mempelajari pola kelas Sentimen secara lebih menyeluruh hingga tidak ada satu pun data Sentimen yang terlewat pada data uji. Aspek Aksesibilitas menunjukkan peningkatan yang paling signifikan dibanding skenario sebelumnya dengan *F1-score* naik dari 0,12 pada skenario 3 menjadi 0,18 pada skenario 4, didukung oleh peningkatan *recall* dari 0,06 menjadi 0,10. Meskipun angka absolut masih rendah, lompatan ini mencerminkan bahwa kombinasi 80% data latih dan *n_estimators* 200 memberikan dampak yang lebih nyata dibanding variasi parameter sebelumnya pada aspek dengan data minoritas yang sangat terbatas. Hal menarik pada skenario 4 adalah presisi Aksesibilitas kembali mencapai 1,00 sama seperti skenario 3 yang mengindikasikan bahwa setiap prediksi Sentimen yang dibuat model selalu tepat, namun model masih sangat konservatif dalam membuat prediksi tersebut. Aspek Amenitas memperoleh *F1-score* sebesar 0,70, sama dengan skenario 2 dan sedikit lebih tinggi dari skenario 3, dengan *recall* 0,59 yang merupakan nilai tertinggi Amenitas sepanjang keempat skenario. Secara keseluruhan tahap Klasifikasi Aspek pada skenario 4 memperoleh rata-rata *F1-score* sebesar 0,61, merupakan nilai tertinggi di antara keempat skenario.

Pada tahap Klasifikasi Sentimen, skenario 4 menunjukkan peningkatan yang lebih terlihat dibanding tiga skenario sebelumnya. Aspek Atraksi kembali memperoleh *F1-score* 0,97 dengan *recall* sempurna 1,00, konsisten di seluruh skenario. Aspek Aksesibilitas memperoleh *F1-score* 0,91, kembali ke level skenario 1 dan 2 setelah sempat turun ke 0,90 pada skenario 3, dengan presisi yang meningkat dari 0,84 menjadi 0,86. Peningkatan ini sejalan dengan bertambahnya

data latih yang memungkinkan model mempelajari pola polaritas Aksesibilitas secara lebih baik. Aspek Amenitas mencatat peningkatan yang paling signifikan pada tahap ini dengan $F1$ -score 0,93, merupakan nilai tertinggi Amenitas di seluruh skenario, naik dari 0,92 pada skenario 1, 2, dan 3. Peningkatan ini didukung oleh kombinasi rasio data latih 80% dan $n_estimators$ 200 yang memberikan dampak sinergis pada aspek dengan distribusi data yang lebih berimbang. Secara keseluruhan tahap Klasifikasi Sentimen pada skenario 4 memperoleh rata-rata $F1$ -score sebesar 0,94, merupakan nilai tertinggi di antara keempat skenario dan satu-satunya skenario di mana Klasifikasi Sentimen menunjukkan peningkatan yang terukur.

Skenario 4 menunjukkan bahwa kombinasi rasio data latih 80% dan $n_estimators$ 200 menghasilkan performa terbaik secara keseluruhan dibanding tiga skenario lainnya, dengan rata-rata $F1$ -score Klasifikasi Aspek 0,61 dan Klasifikasi Sentimen 0,94. Meskipun demikian, peningkatan yang terjadi bersifat inkremental dan tidak mampu mengatasi permasalahan mendasar pada aspek Aksesibilitas yang tetap memperoleh $F1$ -score rendah akibat ketidakseimbangan distribusi data yang ekstrem.

4.2.5 Pembahasan Antar Skenario

Bagian ini membahas perbandingan performa keempat skenario secara menyeluruh guna menjawab tujuan penelitian kedua, yaitu menganalisis pengaruh perubahan $n_estimators$ pada algoritma *Random Forest* untuk mendapatkan model dengan performa paling stabil. Perbandingan dilakukan berdasarkan nilai rata-rata akurasi, presisi, *recall*, dan $F1$ -score dari tahap Klasifikasi Aspek dan Klasifikasi

Sentimen pada setiap skenario, sebagaimana dirangkum pada Tabel 4.17 dan Tabel 4.18 berikut.

Tabel 4.17 Rekapitulasi Performa Klasifikasi Aspek 4 Skenario

Skenario	Rasio	N_estimators	Akurasi	Presisi	Recall	F1-Score
1	70:30	100	0,88	0,89	0,53	0,58
2	70:30	200	0,89	0,90	0,55	0,59
3	80:20	100	0,89	0,92	0,54	0,59
4	80:20	200	0,89	0,92	0,56	0,61

Tabel 4.18 Rekapitulasi Performa Klasifikasi Sentimen 4 Skenario

Skenario	Rasio	N_estimators	Akurasi	Presisi	Recall	F1-Score
1	70:30	100	0,89	0,89	0,99	0,93
2	70:30	200	0,89	0,89	0,99	0,93
3	80:20	100	0,89	0,89	0,98	0,93
4	80:20	200	0,90	0,90	0,98	0,94

Berdasarkan Tabel 4.17 dan Tabel 4.18, terdapat beberapa pola yang dapat diidentifikasi dari perbandingan keempat skenario tersebut. Hasil analisis mendalam terkait perbandingan tersebut menunjukkan adanya lima temuan utama yang dijabarkan dibawah ini.

Pertama, pada tahap Klasifikasi Aspek, terdapat kesenjangan yang konsisten antara nilai akurasi dan nilai *F1-score* serta *recall* di seluruh keempat skenario. Rata-rata akurasi sistem berada di kisaran 0,88–0,89, namun rata-rata *F1-score* hanya berkisar antara 0,58 hingga 0,61 dengan *recall* 0,53 – 0,56. Kondisi ini disebabkan oleh ketidakseimbangan distribusi kelas pada data, di mana model cenderung memprediksi kelas mayoritas sehingga akurasi tetap tinggi namun *F1-score* dan *recall* menjadi rendah. Rendahnya rata-rata *F1-score* dan *recall* secara

keseluruhan dipengaruhi oleh aspek Aksesibilitas yang memiliki ketidakseimbangan distribusi kelas minoritas paling banyak, sehingga menyebabkan turun nilai rata-rata dari ketiga aspek secara drastis. Kondisi ini terjadi di seluruh skenario.

Kedua, pada tahap Klasifikasi Aspek, peningkatan $n_estimators$ dari 100 menjadi 200 memberikan dampak positif meskipun nilainya sedikit. Pada rasio 70:30, peningkatan $n_estimators$ dari 100 ke 200 menghasilkan kenaikan $F1-score$ dari 0,58 menjadi 0,59. Pada rasio 80:20, peningkatan yang sama menghasilkan kenaikan $F1-score$ dari 0,59 menjadi 0,61. Pola ini menunjukkan bahwa penambahan jumlah pohon memberikan pengaruh ketika dikombinasikan dengan rasio data latih yang lebih besar, meskipun selisih tetap kecil.

Ketiga, perubahan rasio pembagian data dari 70:30 menjadi 80:20 memberikan dampak yang lebih konsisten dibanding perubahan $n_estimators$ pada tahap Klasifikasi Aspek. Membandingkan skenario dengan $n_estimators$ yang sama skenario 1 vs 3 dan skenario 2 vs 4 terlihat bahwa peningkatan rasio data latih menghasilkan kenaikan presisi lebih terlihat yaitu dari 0,89 menjadi 0,92 pada kedua pasangan tersebut.

Keempat, pada tahap Klasifikasi Sentimen, model menunjukkan stabilitas yang sangat tinggi di seluruh skenario. Tiga dari empat skenario menghasilkan $F1-score$ yang identik sebesar 0,93, sementara skenario 4 hanya sedikit lebih tinggi di angka 0,94. Rentang variasi $F1-score$ yang hanya sebesar 0,01 di antara keempat skenario mengindikasikan bahwa tahap Klasifikasi Sentimen tidak sensitif terhadap

perubahan *n_estimators* maupun rasio pembagian data model sudah mencapai titik saturasi performa pada parameter ini.

Kelima, skenario 4 dengan rasio 80:20 dan *n_estimators* 200 menghasilkan performa terbaik pada kedua tahap, dengan F1-score Klasifikasi Aspek 0,61 dan Klasifikasi Sentimen 0,94. Meskipun demikian, keunggulan skenario 4 dibanding skenario lainnya tidak terlalu besar, karena selisih F1-score Klasifikasi Aspek antar skenario hanya berkisar antara 0,01 hingga 0,03, dan selisih Klasifikasi Sentimen hanya 0,01. Stabilitas performa yang tinggi di seluruh skenario, terutama pada Klasifikasi Sentimen, menunjukkan bahwa algoritma Random Forest dapat menghasilkan model yang konsisten pada berbagai kombinasi parameter yang diuji dalam penelitian ini. Akan tetapi, masalah ketidakseimbangan distribusi data pada tahap Klasifikasi Aspek tetap menjadi faktor pembatas yang tidak bisa diatasi hanya dengan mengubah parameter.

4.3 Integrasi Islam

Integrasi nilai-nilai Islam dalam penelitian ilmiah merupakan wujud keilmuan yang memadukan antara dimensi spiritual dan rasionalitas teknologi. Dalam konteks ini, pengembangan sistem Analisis Sentimen Berbasis Aspek (ABSA) menggunakan algoritma *Random Forest* tidak sekadar dipandang sebagai penelitian terkait kinerja model, melainkan sebagai bentuk ibadah dan tanggung jawab seorang Muslim. Untuk memahami bagaimana teknologi ini berakar pada nilai-nilai ketauhidan dan sosial, integrasi Islam dalam penelitian ini diklasifikasikan ke dalam dua dimensi utama, yaitu hubungan manusia dengan Sang Pencipta dalam optimalisasi potensi akal (*Muamalah Ma'a Allah Subhanahu wa*

Ta'ala), serta kontribusi nyata dan keadilan informasi yang dihasilkan bagi kemaslahatan masyarakat (*Muamalah Ma'a An-Nas*). Kedua dimensi tersebut diuraikan secara rinci pada bagian berikut.

4.3.1 *Muamalah Ma'a Allah Subhanahu wa Ta'ala*

Islam menempatkan ilmu pengetahuan dan penggunaan akal sebagai kewajiban yang mulia. Allah *Subhanahu wa Ta'ala* tidak menciptakan akal semata sebagai pembeda manusia dengan makhluk lain, melainkan sebagai instrumen untuk memahami, menganalisis dan mengolah segala sesuatu secara sistematis serta bertanggung jawab. Anjuran untuk senantiasa menggunakan akal dan mengembangkan ilmu pengetahuan ditegaskan dalam firman Allah *Subhanahu wa Ta'ala* pada QS. Az-Zumar ayat 9:

قُلْ هَلْ يَسْتَوِي الَّذِينَ يَعْلَمُونَ وَالَّذِينَ لَا يَعْلَمُونَ ۚ إِنَّمَا يَتَذَكَّرُ أُولُو الْأَلْبَابِ

"Katakanlah, apakah sama orang-orang yang mengetahui dengan orang-orang yang tidak mengetahui? Sesungguhnya hanya orang-orang yang berakal sehat yang dapat menerima pelajaran." (QS. Az-Zumar [39]: 9)

Dalam Tafsir Ibnu Katsir, ayat ini menjelaskan bahwa perbedaan antara orang yang berilmu dan orang yang tidak berilmu adalah perbedaan yang nyata dan tidak dapat disamakan. Hanya *ulul albab* yaitu orang-orang yang memiliki akal yang mampu menangkap perbedaan tersebut dan mengambil pelajaran darinya (Al-Sheikh, 2004). Ilmu bukan sekadar pengetahuan yang tersimpan, melainkan sesuatu yang diamalkan dan dimanfaatkan untuk menghasilkan pemahaman yang lebih mendalam terhadap realitas.

Dalam konteks penelitian ini, penerapan algoritma Random Forest pada sistem Analisis Sentimen Berbasis Aspek (ABSA) merupakan wujud nyata dari pengamalan perintah menggunakan akal dan ilmu pengetahuan. Random Forest bekerja dengan membangun sejumlah pohon keputusan yang masing-masing menganalisis fitur-fitur teks secara independen, kemudian mengintegrasikan seluruh hasil analisis tersebut menjadi sebuah keputusan klasifikasi yang lebih akurat dan stabil. Mekanisme ini mencerminkan prinsip *ulul albab* tidak tergesa-gesa dalam mengambil kesimpulan, melainkan mempertimbangkan berbagai sudut pandang secara menyeluruh sebelum menetapkan hasil akhir. Lebih jauh, pendekatan ABSA yang digunakan dalam penelitian ini tidak sekadar mengklasifikasikan sentimen secara umum, melainkan mengurai opini wisatawan ke dalam komponen 3A yang spesifik sehingga pemahaman yang dihasilkan lebih mendalam dan terperinci. Upaya untuk memahami sesuatu secara mendalam dan tidak berhenti pada pemahaman yang dangkal inilah yang sejatinya dimaksud dengan menggunakan akal secara penuh sebagaimana yang diperintahkan Allah *Subhanahu wa Ta'ala* dalam ayat tersebut. Hasil penelitian yang menunjukkan rata-rata F1-score Klasifikasi Sentimen mencapai 0,93 hingga 0,94 di seluruh skenario pengujian merupakan bukti bahwa pemanfaatan ilmu pengetahuan dan teknologi secara sungguh-sungguh dapat menghasilkan pemahaman yang objektif dan terukur terhadap data yang sebelumnya tidak terstruktur.

4.3.2 Muamalah Ma'a An-Nas

Islam tidak hanya mengatur hubungan manusia dengan Allah, tetapi juga menekankan pentingnya memberikan manfaat kepada sesama sebagai bagian dari

tanggung jawab sosial seorang Muslim. Nabi Muhammad *Shallallahu 'Alaihi wa Sallam* bersabda:

خَيْرُ النَّاسِ أَنْفَعُهُمْ لِلنَّاسِ

"Sebaik-baik manusia adalah yang paling bermanfaat bagi manusia lain." (HR. Ahmad, Shahih al-Jami' no 3289)

Hadis ini menegaskan bahwa nilai seorang Muslim tidak hanya diukur dari ibadah ritualnya, melainkan juga dari seberapa besar kontribusi dan manfaat yang diberikan kepada orang lain (Muhammad Nashiruddin AI Albani, 2008). Sejalan dengan itu, Allah *Subhanahu wa Ta'ala* juga memerintahkan manusia untuk berlaku adil dan berbuat kebajikan dalam setiap aspek kehidupan, sebagaimana firman-Nya dalam QS. An-Nahl ayat 90:

إِنَّ اللَّهَ يَأْمُرُ بِالْعَدْلِ وَالْإِحْسَانِ وَإِيتَايِ ذِي الْقُرْبَىٰ وَيَنْهَىٰ عَنِ الْفَحْشَاءِ وَالْمُنْكَرِ وَالْبَغْيِ ۚ يَعِظُكُمْ لَعَلَّكُمْ تَذَكَّرُونَ

"Sesungguhnya Allah menyuruh berlaku adil, berbuat kebajikan, dan memberikan bantuan kepada kerabat. Dia melarang perbuatan keji, kemungkaran, dan permusuhan. Dia memberi pengajaran kepadamu agar kamu dapat mengambil pelajaran." (QS. An-Nahl [16]: 90)

Dalam Tafsir Ibnu Katsir, ayat ini dijelaskan sebagai perintah komprehensif untuk menempatkan segala sesuatu pada proporsinya yang benar dan memberikan hak kepada setiap pihak secara adil (Al-Sheikh, 2004). Prinsip keadilan ini tercermin langsung dalam metodologi penelitian: anotasi data dilakukan secara otomatis menggunakan LLM dengan validasi manual oleh pakar bahasa Indonesia, menghasilkan tingkat kesesuaian 82% yang memastikan setiap ulasan diberi label

secara konsisten dan objektif. Pendekatan ABSA yang digunakan pun menjamin bahwa opini wisatawan tidak digeneralisasi, melainkan diklasifikasikan sesuai tempatnya ke dalam aspek seperti Atraksi, Aksesibilitas, maupun Amenitas yang telah disesuaikan dengan muatan opini sebenarnya dalam setiap ulasan..

Penelitian ini memberikan manfaat bagi dua pihak dalam ruang lingkup penelitian. Pertama, bagi pengembangan ilmu pengetahuan. Penelitian ini memberikan kontribusi empiris dalam kajian *text mining* dan analisis sentimen, khususnya pada penerapan ABSA di domain pariwisata berbahasa Indonesia. Referensi metodologis yang dihasilkan dapat diadopsi dan dikembangkan oleh penelitian selanjutnya untuk memperkaya khazanah keilmuan di bidang ini sebuah bentuk kebajikan ilmiah yang manfaatnya melampaui batas waktu penelitian itu sendiri. Kedua, bagi pengelola destinasi wisata khususnya Jatim Park Group, hasil penelitian ini dapat menjadi dasar dalam mempertimbangkan penggunaan sistem analisis sentimen otomatis untuk memantau opini wisatawan secara lebih efisien. Ketersediaan informasi yang objektif dan terklasifikasi berdasarkan aspek 3A merupakan wujud dari prinsip keadilan dalam ayat tersebut bahwa setiap opini wisatawan didudukkan pada aspeknya yang tepat sesuai tempatnya tanpa generalisasi yang dapat merugikan salah satu pihak.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil penelitian terkait implementasi algoritma *Random Forest* dalam mengklasifikasikan sentimen ulasan wisata Jatim Park Group berbasis aspek 3A serta menganalisis pengaruh perubahan $n_estimators$ terhadap performa model yang telah dilakukan pada keempat skenario, diperoleh kesimpulan sebagai berikut:

1. Performa *Random Forest* dalam Klasifikasi Sentimen Berbasis Aspek

Algoritma *Random Forest* mampu mengklasifikasikan sentimen ulasan wisata Jatim Park Group berbasis aspek 3A. Hasil dari 4 skenario pada penelitian ini menunjukkan bahwa, skenario ke 4 merupakan skenario terbaik pada Klasifikasi Aspek dan Klasifikasi Sentimen dengan pembagian data 80:20 dan $n_estimators$ 200. Berdasarkan skenario ke 4 pada tahap klasifikasi aspek, aspek Atraksi memperoleh performa tertinggi dengan *F1-score* 0,96, diikuti Amenitas dengan *F1-score* 0,70, sementara Aksesibilitas memperoleh *F1-score* terendah sebesar 0,18. Rendahnya performa aspek Aksesibilitas disebabkan karena pengunjung lebih banyak mengulas aspek Atraksi dibandingkan Aksesibilitas, sehingga terjadi ketidakseimbangan distribusi data yang merupakan karakteristik yang melekat pada data ulasan wisata. Sedangkan pada Klasifikasi Sentimen, model menunjukkan performa yang tinggi dengan *F1-score* rata-rata 0,94, di mana aspek Atraksi memperoleh *F1-score* tertinggi sebesar 0,97, diikuti Amenitas 0,93, dan Aksesibilitas 0,91.

2. Pengaruh $n_estimators$ terhadap Performa Klasifikasi

Peningkatan $n_estimators$ dari 100 menjadi 200 memberikan dampak positif namun tidak terlalu besar pada kedua tahap klasifikasi. Pada tahap Klasifikasi Aspek, penambahan jumlah pohon menghasilkan kenaikan rata-rata *F1-score* sebesar 0,01 hingga 0,02, sementara perubahan rasio data latih dari 70:30 menjadi 80:20 memberikan dampak yang lebih terlihat dengan peningkatan presisi dari 0,89 menjadi 0,92. Pada tahap Klasifikasi Sentimen, model menunjukkan nilai yang tinggi yaitu 0,93 dengan rentang nilai *F1-score* hanya sebesar 0,01 di antara keempat skenario, mengindikasikan bahwa model telah mencapai batas nilai tertinggi sehingga tidak sensitif terhadap perubahan parameter yang diuji. Skenario 4 dengan rasio pembagian data 80:20 dan $n_estimators$ 200 menghasilkan performa terbaik secara keseluruhan dengan rata-rata *F1-score* Klasifikasi Aspek 0,61 dan Klasifikasi Sentimen 0,94, sehingga menjadi pengaturan parameter yang paling direkomendasikan dalam penelitian ini.

Dengan demikian, kedua tujuan penelitian telah tercapai melalui pengujian keempat skenario yang dilakukan dalam penelitian ini. Algoritma *Random Forest* terbukti dapat diterapkan dalam klasifikasi sentimen berbasis aspek pada ulasan wisata, meskipun performa model tetap dipengaruhi oleh karakteristik distribusi data yang tidak seimbang. Variasi parameter $n_estimators$ dan rasio pembagian data memberikan dampak yang terukur namun tidak terlalu besar, sehingga pengaturan parameter yang tepat perlu disesuaikan dengan karakteristik data pada domain penelitian yang serupa.

5.2 Saran

Berdasarkan hasil dan keterbatasan yang ditemukan dalam penelitian ini, terdapat beberapa saran untuk pengembangan penelitian selanjutnya.

1. Optimasi penanganan ketidakseimbangan data.

Penelitian selanjutnya disarankan untuk menerapkan teknik *oversampling* seperti *SMOTE* (*Synthetic Minority Oversampling Technique*) sebagai upaya optimasi model agar lebih sensitif dalam mendeteksi aspek dengan jumlah data minoritas. Penerapan SMOTE diarahkan pada level pelatihan model, bukan pada modifikasi data anotasi asli, sehingga integritas representasi opini pengunjung tetap terjaga.

2. Perluasan metode ekstraksi fitur.

Penelitian ini menggunakan TF-IDF sebagai metode representasi fitur. Penelitian berikutnya dapat mengeksplorasi pendekatan berbasis *word embedding* seperti *Word2Vec*, *FastText*, atau model *transformer* seperti IndoBERT untuk menangkap makna semantik yang lebih kaya dari teks ulasan berbahasa Indonesia.

3. Perluasan komparasi algoritma.

Penelitian ini hanya menggunakan *Random Forest* sebagai metode klasifikasi tunggal. Studi lanjutan dapat membandingkan performa *Random Forest* dengan algoritma lain seperti *Support Vector Machine* (SVM), *Gradient Boosting*, atau model *deep learning* untuk mendapatkan referensi perbandingan yang lebih komprehensif.

4. Perluasan cakupan data dan aspek.

Dataset dalam penelitian ini dibatasi pada ulasan Google Maps dan TripAdvisor periode 2020–2025 dengan kerangka aspek 3A. Penelitian selanjutnya dapat memperluas sumber data ke platform lain.

DAFTAR PUSTAKA

- Al-Sheikh, A. Bin M. Bin A. Bin I. (2004). Tafsir Ibnu Katsir Jilid 7 (D. M. Yusuf Harun, Farid Okbah (Ed.)). Pustaka Imam Asy-Syafi'i.
- Amalia Permadi, V. (2020). Analisis Sentimen Menggunakan Algoritma Naive Bayes Terhadap Review Restoran Di Singapura. *Jurnal Buana Informatika*, 11(2), 141–151.
- Andre Setiawan, F. (2025). Kunjungan Wisata 2025 Di Kota Batu Turun Jadi 9,7 Juta. *Radar Batu Jawa Pos*. <https://Radarbatu.Jawapos.Com/Kota-Batu/2327088469/Kunjungan-Wisata-2025-Di-Kota-Batu-Turun-Jadi-97-Juta>
- Angelina, S. J., Bijaksana, A., Negara, P., & Muhardi, H. (2023). Analisis Pengaruh Penerapan Stopword Removal Pada Performa Klasifikasi Sentimen Tweet Bahasa Indonesia. *Juara (Jurnal Aplikasi Dan Riset Informatika)*, 02(1), 165–173. <https://Doi.Org/10.26418/Juara.V2i1.69680>
- Apriani, A., Wahdiniawati, S. A., Saputri, I. P., & Randyantini, V. (2025). Peran Citra Destinasi, Aksesibilitas, Dan Kepuasan Wisatawan Dalam Meningkatkan Niat Kunjungan Ulang Wisatawan Ke Yogyakarta. *Jurnal Bisnis Dan Manajemen*, 5(1), 44–59.
- Ariansyah, K., Prawiro, J., & Sanjaya, R. (2025). Pengaruh Ulasan Online Terhadap Keputusan Wisatawan Dalam Memilih Hotel. *Jurnal Pariwisata Dan Perhotelan*, 2(2), 8. <https://Doi.Org/10.47134/Pjpp.V2i2.3559>
- Arifiyanti, A. A., Pandji, M. F., & Utomo, B. (2022). Analisis Sentimen Ulasan Pengunjung Objek Wisata Gunung Bromo Pada Situs Tripadvisor. *Explore: Jurnal Sistem Informasi Dan Telematika*, 13(1), 32. <https://Doi.Org/10.36448/Jsit.V13i1.2539>
- Asyiah, N., & Aktavia, W. (2025). Penerapan Bertopic Dan Analisis Sentimen Leksikal Pada Ulasan Relevan Di Google Maps Mengenai Universitas Pamulang. *Journal Of Information System Reasearch (Josh)*, 6(4), 2129–2138. <https://Doi.Org/10.47065/Josh.V6i4.8007>
- Aufar, A. F., Mochamad Alfian Rosid, Eviyanti, A., & Astutik, I. R. I. (2023). Optimizing Text Preprocessing For Accurate Sentiment Analysis On E-Wallet Reviews. *Jicte (Journal Of Information And Computer Technology Education)*, 7(2), 42–50. <https://Doi.Org/10.21070/Jicte.V7i2.1650>
- Baskoro, B. B., Susanto, I., & Khomsah, S. (2021). Analisis Sentimen Pelanggan Hotel Di Purwokerto Menggunakan Metode Random Forest Dan Tf-Idf (Studi Kasus: Ulasan Pelanggan Pada Situs Tripadvisor). *Journal Of Informatics, Information System, Software Engineering And Applications*, 3(2), 21–29.

<https://doi.org/10.20895/inista.v3i2>

- Deviyanto, A., & Wahyudi, M. D. R. (2018). Penerapan Analisis Sentimen Pada Pengguna Twitter Menggunakan Metode K-Nearest Neighbor. *Jiska (Jurnal Informatika Sunan Kalijaga)*, 3(1), 1. <https://doi.org/10.14421/jiska.2078.31-01>
- Duei Putri, D., Nama, G. F., & Sulistiono, W. E. (2022). Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (Dpr) Pada Twitter Menggunakan Metode Naive Bayes Classifier. *Jurnal Informatika Dan Teknik Elektro Terapan*, 10(1), 34–40. <https://doi.org/10.23960/jitet.v10i1.2262>
- Jatim Park Group. (2024). Top! Jatim Park Group Raih Penghargaan Kontribusi Pariwisata Disparta Kota Batu. Jatim Park Group Official Website. <https://jtp.id/top-jatim-park-group-raih-penghargaan-kontribusi-pariwisata-disparta-kota-batu/index.html>
- Kecamatan Batu. (N.D.). Mengembangkan Batu Sebagai Kota Wisata Dengan Beragam Atraksi Menarik. Kecamatan Batu Kota Wisata Batu. Retrieved February 14, 2026, From <https://kecamatanbatu.com/mengembangkan-batu-sebagai-kota-wisata-dengan-beragam-atraksi-menarik/>
- Lajnah Pentashihan Mushaf Al-Qur'an. (2019). Al-Qur'an Dan Terjemahannya: Edisi Penyempurnaan 2019. Kementerian Agama RI.
- Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. (2022). Analisis Sentimen Ulasan Aplikasi Dana Dengan Metode Random Forest. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 6(9), 4305–4313.
- Leylita Novita Rossadi, & Endang Widayati. (2024). Pengaruh Aksesibilitas, Amenitas, Dan Atraksi Wisata Terhadap Minat Kunjungan Wisatawan Ke Wahana Air Balong Waterpark Bantul Daerah Istimewa Yogyakarta. *Journal Of Tourism And Economic*, 1(2), 109–116. <https://doi.org/10.36594/jtec/cwkvga87>
- Manullang, O., Prianto, C., & Harani, N. H. (2023). Analisis Sentimen Untuk Memprediksi Hasil Calon Pemilu Presiden Menggunakan Lexicon Based Dan Random Forest. *Jurnal Ilmiah Informatika*, 11(02), 159–169. <https://doi.org/10.33884/jif.v11i02.7987>
- Maulidiya Indah Sari, Masrurotun Mufidah, Sarpini, Y. F. I. (2023). Kontribusi Sektor Pariwisata Terhadap Perekonomian Kebumen. *Jurnal Masharif Al-Syariah: Jurnal Ekonomi Dan Perbankan Syariah*, 8(2), 1177–1217.
- Morama, H. C., Ratnawati, D. E., & Arwani, I. (2022). Analisis Sentimen Berbasis Aspek Terhadap Ulasan Hotel Tentrem Yogyakarta Menggunakan Algoritma

Random Forest Classifier. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 6(4), 1702–1708.

Muhammad Nashiruddin Ai Albani. (2008). *Shahih Ai Jami' Ash-Shaghir Jilid 2*. Pustaka Azzam.

Pontiki, M., Galanis, D., Pavlopoulos, J., Papageorgiou, H., Androutsopoulos, I., & Manandhar, S. (2014). Semeval-2014 Task 4: Aspect Based Sentiment Analysis. *Proceedings Of The 8th International Workshop On Semantic Evaluation (Semeval 2014)*, Semeval, 27–35. <https://doi.org/10.3115/V1/S14-2002>

Pradana, A. (2025). 778.235 Wisatawan Kunjungi Kota Batu Selama Libur Lebaran 2025. *Antara News*. <https://www.antaraneews.com/berita/4765013/778235-Wisatawan-Kunjungi-Kota-Batu-Selama-Libur-Lebaran-2025>

Rofiqi, M. A., Fauzan, A. C., Agustin, A. P., & Saputra, A. A. (2019). Implementasi Term-Frequency Inverse Document Frequency (Tf-Idf) Untuk Mencari Relevansi Dokumen Berdasarkan Query. *Ilkomnika: Journal Of Computer Science And Applied Informatics*, 1(2), 58–64. <https://doi.org/10.28926/Ilkomnika.V1i2.18>

Safari, A., Edison, E., & Gita, V. (2025). Pengaruh Daya Tarik Wisata Dan Aksesibilitas Terhadap Minat Berkunjung Di Kampung Cakrawala Kabupaten Sukabumi. *Manajemen Dan Pariwisata*, 4(1), 135–156. <https://doi.org/10.32659/Jmp.V4i1.416>

Savila, A. G., Supriyono, & Melani, R. I. (2026). Abstractive Summarization Of Indonesian Islamic Stories Using Long Short-Term Memory (Lstm). *Jurnal Teknik Informatika (Jutif)*, 7(1), 158–168.

Septiani, D., & Isabela, I. (2022). Analisis Term Frequency Inverse Document Frequency (Tf-Idf) Dalam Temu Kembali Informasi Pada Dokumen Teks. *Sintesia: Jurnal Sistem Dan Teknologi Informasi Indonesia*, 1(2), 81–88.

Seran, M. Y., Hutagalung, S., Rudiyanto, R., Sandrio, L., & Rostini, I. A. (2023). Analisis Konsep 3a (Atraksi, Amenitas, Akseibilitas) Dalam Perencanaan Pengembangan Pariwisata Berbasis Masyarakat (Studi Kasus: Desa Umatoos, Kabupaten Malaka). *Jptm: Jurnal Penelitian Terapan Mahasiswa*, 1(1), 27–42. <https://journal.poltekkelbajo.ac.id/index.php/jptm/article/view/18>

Stange, J., Brown, D., Hilbruner, R., & Hawkins, D. E. (2010). Tourism Destination Management Achieving Sustainable And Competitive Results. In *Tourism Destination Management*.

Supatmana, R., & Suwarti. (2022). Pengembangan Daya Tarik Wisata Alam Dan

Buatan Berbasis. Jurnal Ekonomi, Manajemen Pariwisata Dan Perhotelan, 1(1).
[Http://Ejournal.Stie-Trianandra.Ac.Id/Index.Php/Jempper](http://ejournal.stie-trianandra.ac.id/index.php/jempper)

Supriyono, Prasetya, A., Kurniawan, F., & Suyono. (2026). Enhancing Stand-Up Comedy Summarization Using Cohesion And Coherence With Deep Learning. *International Journal On Informatics Visualization*, 10(January), 57–65.

Supriyono, Suyono, Wibawa, A. P., & Kurniawan, F. (2024). Analyzing Audience Sentiments In Digital Comedy : A Study Of Youtube Comments Using Lstm Models. *Journal Of Applied Data Sciences*, 5(4), 1877–1889.

Sutriawan, Muljono, Khairunnisa, Alamin, Z., Lorosae, T. A., & Ramadhan, S. (2024). Improving Performance Sentiment Movie Review Classification Using Hybrid Feature Tf-idf, N-Gram, Information Gain And Support Vector Machine. *Mathematical Modelling Of Engineering Problems*, 11(2), 375–384. <https://doi.org/10.18280/mmep.110209>

Sutriawan, Siti Mutmainnah, Teguh Ansyor Lorosae, & Sahrul Ramadhan. (2025). Model Text Embedding Dan Tf-Idf+Ngram Untuk Meningkatkan Kinerja Algoritma Binary Classifier Pada Klasifikasi Sms Palsu. *Jurnal Sistem Informasi Triguna Dharma (Jursi Tgd)*, 4(1), 55–64. <https://doi.org/10.53513/jursi.v4i1.10582>

Al-Sheikh, A. Bin M. Bin A. Bin I. (2004). Tafsir Ibnu Katsir Jilid 7 (D. M. Yusuf Harun, Farid Okbah (Ed.)). Pustaka Imam Asy-Syafi'i.

Amalia Permadi, V. (2020). Analisis Sentimen Menggunakan Algoritma Naive Bayes Terhadap Review Restoran Di Singapura. *Jurnal Buana Informatika*, 11(2), 141–151.

Andre Setiawan, F. (2025). Kunjungan Wisata 2025 Di Kota Batu Turun Jadi 9,7 Juta. *Radar Batu Jawa Pos*. <https://radarbatu.jawapos.com/kota-batu/2327088469/kunjungan-wisata-2025-di-kota-batu-turun-jadi-97-juta>

Angelina, S. J., Bijaksana, A., Negara, P., & Muhandi, H. (2023). Analisis Pengaruh Penerapan Stopword Removal Pada Performa Klasifikasi Sentimen Tweet Bahasa Indonesia. *Juara (Jurnal Aplikasi Dan Riset Informatika)*, 02(1), 165–173. <https://doi.org/10.26418/juara.v2i1.69680>

Apriani, A., Wahdiniawati, S. A., Saputri, I. P., & Randyantini, V. (2025). Peran Citra Destinasi, Aksesibilitas, Dan Kepuasan Wisatawan Dalam Meningkatkan Niat Kunjungan Ulang Wisatawan Ke Yogyakarta. *Jurnal Bisnis Dan Manajemen*, 5(1), 44–59.

Ariansyah, K., Prawiro, J., & Sanjaya, R. (2025). Pengaruh Ulasan Online Terhadap

Keputusan Wisatawan Dalam Memilih Hotel. *Jurnal Pariwisata Dan Perhotelan*, 2(2), 8. <https://doi.org/10.47134/Pjpp.V2i2.3559>

Arifiyanti, A. A., Pandji, M. F., & Utomo, B. (2022). Analisis Sentimen Ulasan Pengunjung Objek Wisata Gunung Bromo Pada Situs Tripadvisor. *Explore: Jurnal Sistem Informasi Dan Telematika*, 13(1), 32. <https://doi.org/10.36448/Jsit.V13i1.2539>

Asyiah, N., & Aktavia, W. (2025). Penerapan Bertopic Dan Analisis Sentimen Leksikal Pada Ulasan Relevan Di Google Maps Mengenai Universitas Pamulang. *Journal Of Information System Reasearch (Josh)*, 6(4), 2129–2138. <https://doi.org/10.47065/Josh.V6i4.8007>

Aufar, A. F., Mochamad Alfian Rosid, Eviyanti, A., & Astutik, I. R. I. (2023). Optimizing Text Preprocessing For Accurate Sentiment Analysis On E-Wallet Reviews. *Jicte (Journal Of Information And Computer Technology Education)*, 7(2), 42–50. <https://doi.org/10.21070/Jicte.V7i2.1650>

Baskoro, B. B., Susanto, I., & Khomsah, S. (2021). Analisis Sentimen Pelanggan Hotel Di Purwokerto Menggunakan Metode Random Forest Dan Tf-Idf (Studi Kasus: Ulasan Pelanggan Pada Situs Tripadvisor). *Journal Of Informatics, Information System, Software Engineering And Applications*, 3(2), 21–29. <https://doi.org/10.20895/Inista.V3i2>

Deviyanto, A., & Wahyudi, M. D. R. (2018). Penerapan Analisis Sentimen Pada Pengguna Twitter Menggunakan Metode K-Nearest Neighbor. *Jiska (Jurnal Informatika Sunan Kalijaga)*, 3(1), 1. <https://doi.org/10.14421/Jiska.2018.31-01>

Duei Putri, D., Nama, G. F., & Sulistiono, W. E. (2022). Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (Dpr) Pada Twitter Menggunakan Metode Naive Bayes Classifier. *Jurnal Informatika Dan Teknik Elektro Terapan*, 10(1), 34–40. <https://doi.org/10.23960/Jitet.V10i1.2262>

Jatim Park Group. (2024). Top! Jatim Park Group Raih Penghargaan Kontribusi Pariwisata Disparta Kota Batu. *Jatim Park Group Official Website*. <https://jtp.id/top-jatim-park-group-raih-penghargaan-kontribusi-pariwisata-disparta-kota-batu/index.html>

Kecamatan Batu. (N.D.). Mengembangkan Batu Sebagai Kota Wisata Dengan Beragam Atraksi Menarik. *Kecamatan Batu Kota Wisata Batu*. Retrieved February 14, 2026, From <https://kecamatanbatu.com/mengembangkan-batu-sebagai-kota-wisata-dengan-beragam-atraksi-menarik/>

Lajnah Pentashihan Mushaf Al-Qur'an. (2019). *Al-Qur'an Dan Terjemahannya: Edisi Penyempurnaan 2019*. Kementerian Agama RI.

- Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. (2022). Analisis Sentimen Ulasan Aplikasi Dana Dengan Metode Random Forest. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 6(9), 4305–4313.
- Leylita Novita Rossadi, & Endang Widayati. (2024). Pengaruh Aksesibilitas, Amenitas, Dan Atraksi Wisata Terhadap Minat Kunjungan Wisatawan Ke Wahana Air Balong Waterpark Bantul Daerah Istimewa Yogyakarta. *Journal Of Tourism And Economic*, 1(2), 109–116. <https://doi.org/10.36594/Jtec/Cwkvga87>
- Manullang, O., Prianto, C., & Harani, N. H. (2023). Analisis Sentimen Untuk Memprediksi Hasil Calon Pemilu Presiden Menggunakan Lexicon Based Dan Random Forest. *Jurnal Ilmiah Informatika*, 11(02), 159–169. <https://doi.org/10.33884/Jif.V11i02.7987>
- Maulidiya Indah Sari, Masrurotun Mufidah, Sarpini, Y. F. I. (2023). Kontribusi Sektor Pariwisata Terhadap Perekonomian Kebumen. *Jurnal Masharif Al-Syariah: Jurnal Ekonomi Dan Perbankan Syariah*, 8(2), 1177–1217.
- Morama, H. C., Ratnawati, D. E., & Arwani, I. (2022). Analisis Sentimen Berbasis Aspek Terhadap Ulasan Hotel Tentrem Yogyakarta Menggunakan Algoritma Random Forest Classifier. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 6(4), 1702–1708.
- Muhammad Nashiruddin Ai Albani. (2008). *Shahih Ai Jami' Ash-Shaghir Jilid 2*. Pustaka Azzam.
- Pontiki, M., Galanis, D., Pavlopoulos, J., Papageorgiou, H., Androutsopoulos, I., & Manandhar, S. (2014). Semeval-2014 Task 4: Aspect Based Sentiment Analysis. *Proceedings Of The 8th International Workshop On Semantic Evaluation (Semeval 2014)*, Semeval, 27–35. <https://doi.org/10.3115/V1/S14-2002>
- Pradana, A. (2025). 778.235 Wisatawan Kunjungi Kota Batu Selama Libur Lebaran 2025. *Antara News*. <https://www.antarane.ws.com/Berita/4765013/778235-Wisatawan-Kunjungi-Kota-Batu-Selama-Libur-Lebaran-2025>
- Rofiqi, M. A., Fauzan, A. C., Agustin, A. P., & Saputra, A. A. (2019). Implementasi Term-Frequency Inverse Document Frequency (Tf-Idf) Untuk Mencari Relevansi Dokumen Berdasarkan Query. *Ilkomnika: Journal Of Computer Science And Applied Informatics*, 1(2), 58–64. <https://doi.org/10.28926/Ilkomnika.V1i2.18>
- Safari, A., Edison, E., & Gita, V. (2025). Pengaruh Daya Tarik Wisata Dan Aksesibilitas Terhadap Minat Berkunjung Di Kampung Cakrawala Kabupaten

Sukabumi. *Manajemen Dan Pariwisata*, 4(1), 135–156.
<https://doi.org/10.32659/jmp.v4i1.416>

Savila, A. G., Supriyono, & Melani, R. I. (2026). Abstractive Summarization Of Indonesian Islamic Stories Using Long Short-Term Memory (Lstm). *Jurnal Teknik Informatika (Jutif)*, 7(1), 158–168.

Septiani, D., & Isabela, I. (2022). Analisis Term Frequency Inverse Document Frequency (Tf-Idf) Dalam Temu Kembali Informasi Pada Dokumen Teks. *Sintesia: Jurnal Sistem Dan Teknologi Informasi Indonesia*, 1(2), 81–88.

Seran, M. Y., Hutagalung, S., Rudiyanto, R., Sandrio, L., & Rostini, I. A. (2023). Analisis Konsep 3a (Atraksi, Amenitas, Akseibilitas) Dalam Perencanaan Pengembangan Pariwisata Berbasis Masyarakat (Studi Kasus: Desa Umatoos, Kabupaten Malaka). *Jptm: Jurnal Penelitian Terapan Mahasiswa*, 1(1), 27–42.
<https://journal.poltekkelbajo.ac.id/index.php/jptm/article/view/18>

Stange, J., Brown, D., Hilbruner, R., & Hawkins, D. E. (2010). Tourism Destination Management Achieving Sustainable And Competitive Results. In *Tourism Destination Management*.

Supatmana, R., & Suwarti. (2022). Pengembangan Daya Tarik Wisata Alam Dan Buatan Berbasis. *Jurnal Ekonomi, Manajemen Pariwisata Dan Perhotelan*, 1(1).
<http://ejurnal.stie-trianandra.ac.id/index.php/jempper>

Supriyono, Prasetya, A., Kurniawan, F., & Suyono. (2026). Enhancing Stand-Up Comedy Summarization Using Cohesion And Coherence With Deep Learning. *International Journal On Informatics Visualization*, 10(January), 57–65.

Supriyono, Suyono, Wibawa, A. P., & Kurniawan, F. (2024). Analyzing Audience Sentiments In Digital Comedy : A Study Of Youtube Comments Using Lstm Models. *Journal Of Applied Data Sciences*, 5(4), 1877–1889.

Sutriawan, Muljono, Khairunnisa, Alamin, Z., Lorosae, T. A., & Ramadhan, S. (2024). Improving Performance Sentiment Movie Review Classification Using Hybrid Feature Tfidf, N-Gram, Information Gain And Support Vector Machine. *Mathematical Modelling Of Engineering Problems*, 11(2), 375–384.
<https://doi.org/10.18280/mmep.110209>

Sutriawan, Siti Mutmainnah, Teguh Ansyor Lorosae, & Sahrul Ramadhan. (2025). Model Text Embedding Dan Tf-Idf+Ngram Untuk Meningkatkan Kinerja Algoritma Binary Classifier Pada Klasifikasi Sms Palsu. *Jurnal Sistem Informasi Triguna Dharma (Jursi Tgd)*, 4(1), 55–64. <https://doi.org/10.5353/jursi.v4i1.10582>

LAMPIRAN

Lampiran 1 : Hasil Anotasi LLM (Claude)

No	review_teks	atraksi	amenitas	aksesibilitas
1	Bagus. Tempat yang harus dikunjungi apalagi buat keluarga yang punya anak kecil. Ini sangat direkomendasikan karena bisa memberikan edukasi ke anak" tentang binatang. Punya tempat bermainnya juga dan semua wahana bisa dicoba. Untuk tempat makan didalamnya juga ada banyak pilihan. Rest room juga banyak tempatnya. Disarankan kalo mau datang dari pagi aja apalagi pas weekend ato waktu libur karena pasti padat pengunjung. Kalo cuma setengah hari sayang gak bisa cobain semuanya. Harga pakatnya utk sekarang 160K/org sudah untuk semua tempat bisa di kunjungi. Diluar makan sama bbrpa transportasi kayak sepeda kecil buat keliling taman bagi yang cape jalan kaki. Tempatnya edukatif kebersihannya jg cukup baik dan kalo datang kesini kita juga membantu komunitasnya untuk mempertahankan kelangsungan hidup binatang"nya. Smg hewan"nya sehat selalu sama karyawannya jg. Hehe	Positif	None	Positif
2	Seru, paket nya lumayan lengkap main disini. Bisa keliling lihat satwa, lihat pertunjukan satwa, main, dan foto sama satwa dan ke baby zoo.... tapi sayangnya kolam renangya cocok untuk children aja.Laper, jangan khawatir,ada foodcord disana. Puas keliling area satwa hidup, Lanjut ke museum satwa...	Positif	None	Netral
3	Wahananya masih bagus2 spt dulu...jg ada wahana baru.perawatan tempatnya keren..dicat lagi..	Positif	None	None
4	Jatim Park 3 menurutku adalah wisata modern di kota Malang yang sangat kerennn 🤩🤩 Salah satu daya tarik utamanya adalah Dino Park, di dalam Dino Park sangat menyenangkan dan dapat mengetahui replika dinosaurus yang sangat realis.	Positif	None	Positif

No	review_teks	atraksi	amenitas	aksesibilitas
	Tempat ini juga punya wahana indoor seperti Museum Musik Dunia yang lengkap dengan berbagai penjelasan dan sejarahnya, Fun Tech Plaza, dan Ice Age, yang semuanya menyuguhkan pengalaman interaktif dan menyenangkan. Fasilitasnya juga sangat lengkap. HTM cukup sepadan dengan pengalaman yang didapat. Cocok dikunjungi oleh semua usia, terutama anak-anak dan remaja. Saran: datang lebih pagi agar bisa menikmati semua wahana dengan maksimal REKOMENNNN banget deh buat kalian yang mau ke JATIM PARK 3			
5	Beberapa waktu saya berkunjung bersama rombongan study tour dari sekolah saya untuk berwisata di kawasan Kota Batu. Selama berwisata, salah satu destinasi yang dikunjungi adalah Jatim Park 1. Di sini, saya dan rombongan mengeksplor beragam hal. Mulai dari area edukasi, sejarah, budaya, dan kebun binatang. Lalu selanjutnya akan ada area rekreasi berupa wahana yang cukup beragam. Mulai dari untuk bersenang ataupun memacu adrenalin. Selain itu, kawasan kolam renang juga ada dan lengkap. Jatim Park 1 memiliki harga tiket masuk yang cukup lumayan murah, tergantung umur dan tinggi. Saya merekomendasikan untuk wisata keluarga disini.	Positif	None	None
6	Tempat pengenalan berbagai macam satwa dan aneka burung, ada area kolam renang dan taman bermain, cocok untuk anak" dalam pengenalan satwa, tempatnya sejuk dan ada area food court juga	Positif	None	Positif
7	Jatim park tidak henti-hentinya memberikan inovasi tempat hiburan di kota Batu, Malang. Yang terbaru adalah Dino Park di Jatim Park 3. Tempat ini mengusung tema khusus yaitu Dinosaur,	Positif	None	Positif

No	review_teks	atraksi	amenitas	aksesibilitas
	namun tidak hanya semata-mata Dinosaurus yang bagus ditempat ini namun ada juga didalamnya The Legend Star, dimana kita bisa merasa keliling dunia, dengan bertemu para tokoh ternama di seluruh dunia , mirip dengan madame Tussaud tetapi menurut saya lebih puas kesini karena tidak hanya menyajikan patung2 para tokoh tetapi juga berbagai konsep tempat terkenal diseluruh dunia maupun dari cerita dan film-film seperti Diagon Alley tempat Harry Potter belanja buku-buku. Selain itu ada juga berbagai tempat makan yang ada di Dino Mall nya. Ada juga Fun Tech untuk anak-anak, ada juga Magical Circus yang dibuka ditempat ini. Tempatnya cukup luas, dan usahakan berangkat lebih awal karena antusias pengunjung ditempat ini sangat tinggi. Untuk Dinopark saat ini kondisi di siang hari sangat-sangat panas jangan lupa bawa peralatan seperti topi, payung,air minum,kacamata, namun tempat ini juga tutup cukup awal sekitar pkl 16.00. Bagus memang karena teknologi yang digunakan sudah canggih seperti di LN . Saya rasa tidak akan rugi jika pergi ke tempat ini bersama seluruh anggota keluarga. Untuk tiket pertimbangkan sesuai budget dan tujuan anda liburan kemari, karena dalam satu hari tidak akan cukup untuk menyusuri seluruh tempat ini dan mengikuti permainannya.			
8	Selalu suka sama tempat wisata ini,toiletnya bersih Dan lengkap untuk disabilitas Dan lansia pun ada,bersinergi dengan UMKM dg adanya space UMKM,cuma sedih ga bisa ngejajal semua wahana wisata saking banyaknya,selalu ngangenin	Positif	Positif	Positif
9	Eco Green Park..... Tempat wisata ini adalah salah satu tempat wisata terbaik di Kota Batu - Malang. Sekilas seperti berada dalam Kebun Binatang	Positif	None	None

No	review_teks	atraksi	amenitas	aksesibilitas
	yang dikelilingi penuh dengan Hewan - Hewan. Tetapi jangan salah... Eco Green Park ini jauh lebih bagus daripada Kebun Binatang... Kita bisa menjumpai Hewan - Hewan yang tidak bisa kita temukan di Kebun Binatang biasa. Ada banyak jenis Hewan yang berkeliaran disini... Walaupun banyak yang berkeliaran namun masih tetap berada dalam Zona Area Hewan itu sendiri. Tempat ini, Eco Green Park adalah tempat yang cocok untuk Spot Foto - Foto apalagi jika bersama dengan pasangan kita. Buat saya Eco Green Park ini adalah tempat yang menurut saya tempat wisata Terbaik di Kota Batu - Malang. Eiiitt.... Kita juga bisa berfoto bersama Burung Elang yang tentunya selalu dipandu oleh Pelatih Elang nya. Saat berfoto sang pelatih akan menyuruh kita menggunakan Sarung Tangan khusus sebagai tempat berpijaknya Burung Elang di tangan kita. Jangan takut yach, Guys.. Burung Elangnya jinak kok, meskipun kelihatan Burung Elangnya bertampang agak - agak serem. Bagi yang belum pernah berkunjung ke tempat wisata Eco Green Park ini... Anda - Anda wajib mengunjungi tempat wisata yang satu ini.... Selamat bersenang - senang, Guys..... 👍 👍 👍			
10	Bagus sekali. Koleksi mobil-mobil dan motor-motor kunonya banyak sekali. Benar-benar takjub jika berkunjung ke sini. Yang mengesankan adalah setiap spot berbeda-beda temanya. Jadi berasa ada di China, Amerika, Eropa dalam sehari. Salut kepada pengelola Museum Angkut yang punya konsep luar biasa. Keluar dari area Museum Angkut akan ketemu spot-spot makan dengan berbagai tema negara, korea, china, arab, indonesia, dll.	Positif	None	Positif

No	review_teks	atraksi	amenitas	aksesibilitas
11	Jatimpark 2 / secret zoo, merupakan tempat rekreasi sekaligus edukasi yg cocok untuk anak2 belajar mengenal hewan. Banyak pertunjukan yg diselenggarakan pada jam tertentu. Harga tiket masuk berbanding lurus dengan pelayanan yg diberikan.	Positif	None	None
12	Seru banget, banyak banget wahananya, ditambah lagi wahana yg nggak include sama paket murah-murah lagi Makananya juga harganya wajar dan lumayan enak	Netral	None	Positif
13	Beberapa hari lalu kesana dan harga tiket sekitar Rp 130 ribu, untuk lokasi mudah dijangkau dan sangat rekomendasi untuk liburan keluarga. Untuk air minum & makanan bayi diperbolehkan dibawa masuk. Dan sangat2 puas sekali karena lokasinya sangat nyaman bersih dan aman. See ya next time.	None	Positif	Positif
14	Jatim Park 2 yg berlokasi di kota wisata Batu, Malang sangat cocok untuk rekreasi keluarga apalagi bagi yg memiliki anak kecil, karena disamping tempat bermain juga sebagai sarana pendidikan dan mengenal lingkungan hidup. Lokasi cukup bersih dan terawat dan harga tiket yg terjangkau.	Positif	None	Positif
15	Sebenarnya sudah sangat sering saya dan keluarga berkunjung kesini, karena selain dekat dengan tempat tinggal kami di Surabaya dan Malang juga hawa sejuknya kota Batu. Apalagi di Jatim Park 2 menawarkan berbagai objek menarik untuk berlibur sembari mengedukasi saya dan keluarga, terutama anak-anak. Bagi saya Jatim Park 2 adalah objek wisata yg sangat bagus untuk tujuan liburan keluarga. Apalagi semakin lama management semakin berbenah dengan terus menambah wahana. Sukses terus untuk Jatim Park 2	Positif	None	Positif
16	sesuai ekspektasi banget tempat nya, kemarin aku kesana harga tiket masuk nya per orang 110k itu	Positif	None	Positif

No	review_teks	atraksi	amenitas	aksesibilitas
	udaa termasuk naik wahana sepuasnya. dan untuk makanan dan minumannya juga rekomen banget murah ² g begitu mahal. buat kalian yang mau kesini langsung aja dijamin g akan nyesel kok karna tempat nya melebihi ekspektasi kita 😊			
17	Cocok untuk tempat mempelajari budaya maupun pengetahuan tentang tubuh, bermain, bersenang-senang, dan berbelanja bersama teman atau pun keluarga, karena disana memiliki berbagai jenis mainan dan disediakan pula tempat untuk membeli oleh-oleh.	Positif	None	Positif
18	Tempat wisata yg bagus, cocok u/ Keluarga n edukasi anak2. Dengan memakai boarding pass maskapai tertentu bisa dapat diskon sampai 20%.	Positif	None	None
19	Anakku suka banget diajak kesini 🥰 Tempat edukasi rekomen kalau ke Kota Batu Ada kolam buat bermain air cuma tadi pas kesana kolam nya agak kotor, ada juga wahana permainan nya, cuma ada batasan umur kalau mau naik. Ada wahana berbayar juga seperti flying fox dan sinema 3D Kalau dari tempat parkir motor jalannya agak jauh kalau ke loket. Paling favorit di tempat merak putih 🐦💕	Netral	Negatif	Negatif
20	Gak cukup sehari kalo maen kesini karena luas banget dengan koleksi hewan yg banyak. Selain itu ada wahana bermain dan kolam renang anak. Wahana di dalam mayoritas sudah include di tiket masuk jadi gak perlu bayar lagi. Beberapa wahana ...	Positif	None	None
21	Tempat wisata yang cocok untuk keluarga. Saya mendatangi DINO PARK dan Millennial Glow Garden. Tempat yang sangat indah dan edukatif. Penerapan protokol kesehatan sangat ketat. Tersedia juga penyewaan kursi roda gratis untuk	Positif	None	Positif

No	review_teks	atraksi	amenitas	aksesibilitas
	lansia. Memiliki konsep yang sangat millenial dan keren..			
22	Alhamdulillah JP2 sudah dibuka kembali. Wahana2nya banyak yg baru, apa mungkin krn sy lama gak ke sini lagi ya jd gak tau. Protokol kesehatan jg sangat diperhatikan mulai pintu masuk sampai di dalam, hingga keluar.	Positif	None	Netral
23	Pas kesana rame full mengantri parah,panas dan tempat penuh sampai gk ada tempat buat duduk.semua wahana juga penuh antri ² berdesakan ,kolam anak penuh .. ada beberapa wahana yg rusak tidak bisa dinaiki,, trus juga banyak yg diperbaiki ...	Negatif	None	None
24	Tujuan perjalanan liburan ini menggunakan konsep edukatif yang menggabungkan dunia museum dinosaurus di abad-abad dan dunia legendaris alat musik dan seniman, negara dan pemimpin mereka, film dan seniman, olahraga dan seniman. Pengalaman liburan sangat menyenangkan. Buka dari 11 siang hingga 10 malam pada akhir pekan. Silakan menikmati dunia legendaris museum pertama, maka dunia museum dinosaurus. Ada juga berbagai kuliner, musik live dan merchandise. Hmmm bisa dalam satu hari tidak cukup untuk menikmati keseluruhan, jadi bagikan waktu Anda dengan baik. Nikmati bersama keluarga, kerabat dan teman-teman ...	Positif	None	Netral

Lampiran 2 : Hasil Anotasi Pakar (Yoga Yolanda, M.Pd)

No	review_teks	atraksi	amenitas	aksesibilitas
1	Bagus. Tempat yang harus dikunjungi apalagi buat keluarga yang punya anak kecil. Ini sangat direkomendasikan karena bisa memberikan edukasi ke anak" tentang binatang. Punya tempat bermainnya juga dan semua wahana bisa dicoba.	Positif	Positif	Positif

No	review_teks	atraksi	amenitas	aksesibilitas
	Untuk tempat makan didalamnya juga ada banyak pilihan. Rest room juga banyak tempatnya. Disarankan kalo mau datang dari pagi aja apalagi pas weekend ato waktu libur karena pasti padat pengunjung. Kalo cuma setengah hari sayang gak bisa cobain semuanya. Harga pakatnya utk sekarang 160K/org sudah untuk semua tempat bisa di kunjungi. Diluar makan sama bbrpa transportasi kayak sepeda kecil buat keliling taman bagi yang cape jalan kaki. Tempatnya edukatif kebersihannya jg cukup baik dan kalo datang kesini kita juga membantu komunitasnya untuk mempertahankan kelangsungan hidup binatang"nya. Smg hewan"nya sehat selalu sama karyawannya jg. Hehe			
2	Seru, paket nya lumayan lengkap main disini. Bisa keliling lihat satwa, lihat pertunjukan satwa, main, dan foto sama satwa dan ke baby zoo.... tapi sayangnya kolam renangya cocok untuk children aja.Laper, jangan khawatir,ada foodcord disana. Puas keliling area satwa hidup, Lanjut ke musium satwa...	Positif	None	None
3	Wahanya masih bagus2 spt dulu...jg ada wahana baru.perawatan tempatnya keren..dicat lagi..	Positif	None	None
4	Jatim Park 3 menurutku adalah wisata modern di kota Malang yang sangat kerennn 🥳🥳 Salah satu daya tarik utamanya adalah Dino Park, di dalam Dino Park sangat menyenangkan dan dapat mengetahui replika dinosaurus yang sangat realis. Tempat ini juga punya wahana indoor seperti Museum Musik Dunia yang lengkap dengan berbagai penjelasan dan sejarahnya, Fun Tech Plaza, dan Ice Age, yang semuanya menyuguhkan pengalaman interaktif dan menyenangkan. Fasilitasnya jugaa sangat lengkapp. HTM cukup	Positif	Positif	None

No	review_teks	atraksi	amenitas	aksesibilitas
	sepadan dengan pengalaman yang didapat. Cocok dikunjungi oleh semua usia, terutama anak-anak dan remaja. Saran: datang lebih pagi agar bisa menikmati semua wahana dengan maksimal REKOMENNNN banget deh buat kalian yang mau ke JATIM PARK 3			
5	Beberapa waktu saya berkunjung bersama rombongan study tour dari sekolah saya untuk berwisata di kawasan Kota Batu. Selama berwisata, salah satu destinasi yang dikunjungi adalah Jatim Park 1. Di sini, saya dan rombongan mengeksplor beragam hal. Mulai dari area edukasi, sejarah, budaya, dan kebun binatang. Lalu selanjutnya akan ada area rekreasi berupa wahana yang cukup beragam. Mulai dari untuk bersenang ataupun memacu adrenalin. Selain itu, kawasan kolam renang juga ada dan lengkap. Jatim Park 1 memiliki harga tiket masuk yang cukup lumayan murah, tergantung umur dan tinggi. Saya merekomendasikan untuk wisata keluarga disini.	Positif	None	None
6	Tempat pengenalan berbagai macam satwa dan aneka burung, ada area kolam renang dan taman bermain, cocok untuk anak" dalam pengenalan satwa, tempatnya sejuk dan ada area food court juga	Netral	Positif	None
7	Jatim park tidak henti-hentinya memberikan inovasi tempat hiburan di kota Batu, Malang. Yang terbaru adalah Dino Park di Jatim Park 3. Tempat ini mengusung tema khusus yaitu Dinosaur, namun tidak hanya semata-mata Dinosaur yang bagus ditempat ini namun ada juga didalamnya The Legend Star, dimana kita bisa merasa keliling dunia, dengan bertemu para tokoh ternama di seluruh dunia , mirip dengan madame Tussaud tetapi menurut saya lebih puas kesini karena tidak	Positif	Positif	None

No	review_teks	atraksi	amenitas	aksesibilitas
	hanya menyajikan patung2 para tokoh tetapi juga berbagai konsep tempat terkenal diseluruh dunia maupun dari cerita dan film-film seperti Diagon Alley tempat Harry Potter belanja buku-buku. Selain itu ada juga berbagai tempat makan yang ada di Dino Mall nya. Ada juga Fun Tech untuk anak-anak, ada juga Magical Circus yang dibuka ditempat ini. Tempatnya cukup luas, dan usahakan berangkat lebih awal karena antusias pengunjung ditempat ini sangat tinggi. Untuk Dinopark saat ini kondisi di siang hari sangat-sangat panas jangan lupa bawa peralatan seperti topi, payung,air minum,kacamata, namun tempat ini juga tutup cukup awal sekitar pkl 16.00. Bagus memang karena teknologi yang digunakan sudah canggih seperti di LN . Saya rasa tidak akan rugi jika pergi ke tempat ini bersama seluruh anggota keluarga. Untuk tiket pertimbangkan sesuai budget dan tujuan anda liburan kemari, karena dalam satu hari tidak akan cukup untuk menyusuri seluruh tempat ini dan mengikuti permainannya.			
8	Selalu suka sama tempat wisata ini,toiletnya bersih Dan lengkap untuk disabilitas Dan lansia pun ada,bersinergi dengan UMKM dg adanya space UMKM,cuma sedih ga bisa ngejajal semua wahana wisata saking banyaknya,selalu ngangenin	Positif	Positif	None
9	Eco Green Park..... Tempat wisata ini adalah salah satu tempat wisata terbaik di Kota Batu - Malang. Sekilas seperti berada dalam Kebun Binatang yang dikelilingi penuh dengan Hewan - Hewan. Tetapi jangan salah... Eco Green Park ini jauh lebih bagus daripada Kebun Binatang... Kita bisa menjumpai Hewan - Hewan yang tidak bisa kita temukan di Kebun Binatang biasa. Ada banyak jenis Hewan yang berkeliaran disini... Walaupun	Positif	None	Negatif

No	review_teks	atraksi	amenitas	aksesibilitas
	banyak yang berkeliaran namun masih tetap berada dalam Zona Area Hewan itu sendiri. Tempat ini, Eco Green Park adalah tempat yang cocok untuk Spot Foto - Foto apalagi jika bersama dengan pasangan kita. Buat saya Eco Green Park ini adalah tempat yang menurut saya tempat wisata Terbaik di Kota Batu - Malang. Eiiitt.... Kita juga bisa berfoto bersama Burung Elang yang tentunya selalu dipandu oleh Pelatih Elang nya. Saat berfoto sang pelatih akan menyuruh kita menggunakan Sarung Tangan khusus sebagai tempat berpijaknya Burung Elang di tangan kita. Jangan takut yach, Guys.. Burung Elangnya jinak kok, meskipun kelihatan Burung Elangnya bertampang agak - agak serem. Bagi yang belum pernah berkunjung ke tempat wisata Eco Green Park ini... Anda - Anda wajib mengunjungi tempat wisata yang satu ini.... Selamat bersenang - senang, Guys..... 👍 👍 👍			
10	Bagus sekali. Koleksi mobil-mobil dan motor-motor kunonya banyak sekali. Benar-benar takjub jika berkunjung ke sini. Yang mengesankan adalah setiap spot berbeda-beda temanya. Jadi berasa ada di China, Amerika, Eropa dalam sehari. Salut kepada pengelola Museum Angkut yang punya konsep luar biasa. Keluar dari area Museum Angkut akan ketemu spot-spot makan dengan berbagai tema negara, korea, china, arab, indonesia, dll.	Positif	Positif	None
11	Jatimpark 2 / secret zoo, merupakan tempat rekreasi sekaligus edukasi yg cocok untuk anak2 belajar mengenal hewan. Banyak pertunjukan yg diselenggarakan pada jam tertentu. Harga tiket masuk berbanding lurus dengan pelayanan yg diberikan.	Positif	None	None

No	review_teks	atraksi	amenitas	aksesibilitas
12	Seru banget, banyak banget wahananya, ditambah lagi wahana yg nggak include sama paket murah-murah lagi Makananya juga harganya wajar dan lumayan enak	Positif	Positif	None
13	Beberapa hari lalu kesana dan harga tiket sekitar Rp 130 ribu, untuk lokasi mudah dijangkau dan sangat rekomendasi untuk liburan keluarga. Untuk air minum & makanan bayi diperbolehkan dibawa masuk. Dan sangat2 puas sekali karena lokasinya sangat nyaman bersih dan aman. See ya next time.	None	Positif	Positif
14	Jatim Park 2 yg berlokasi di kota wisata Batu, Malang sangat cocok untuk rekreasi keluarga apalagi bagi yg memiliki anak kecil, karena disamping tempat bermain juga sebagai sarana pendidikan dan mengenal lingkungan hidup. Lokasi cukup bersih dan terawat dan harga tiket yg terjangkau.	None	Positif	Netral
15	Sebenarnya sudah sangat sering saya dan keluarga berkunjung kesini, karena selain dekat dengan tempat tinggal kami di Surabaya dan Malang juga hawa sejuknya kota Batu. Apalagi di Jatim Park 2 menawarkan berbagai objek menarik untuk berlibur sembari mengedukasi saya dan keluarga, terutama anak-anak. Bagi saya Jatim Park 2 adalah objek wisata yg sangat bagus untuk tujuan liburan keluarga. Apalagi semakin lama management semakin berbenah dengan terus menambah wahana. Sukses terus untuk Jatim Park 2	Positif	Positif	Positif
16	sesuai ekspektasi banget tempat nya, kemarin aku kesana harga tiket masuk nya per orang 110k itu udaa termasuk naik wahana sepuasnya. dan untuk makanan dan minumannya juga rekomen banget murah² g begitu mahal. buat kalian yang mau kesini langsung aja dijamin g akan nyesel kok karna tempat nya melebihi ekspektasi kita 😊	Positif	Netral	None

No	review_teks	atraksi	amenitas	aksesibilitas
17	Cocok untuk tempat mempelajari budaya maupun pengetahuan tentang tubuh, bermain, bersenang-senang, dan berbelanja bersama teman atau pun keluarga, karena disana memiliki berbagai jenis mainan dan disediakan pula tempat untuk membeli oleh-oleh.	Positif	Positif	None
18	Tempat wisata yg bagus, cocok u/ Keluarga n edukasi anak2. Dengan memakai boarding pass maskapai tertentu bisa dapat diskon sampai 20%.	Positif	None	None
19	Anakku suka banget diajak kesini 🥰 Tempat edukasi rekomen kalau ke Kota Batu Ada kolam buat bermain air cuma tadi pas kesana kolam nya agak kotor, ada juga wahana permainan nya, cuma ada batasan umur kalau mau naik. Ada wahana berbayar juga seperti flying fox dan sinema 3D Kalau dari tempat parkir motor jalannya agak jauh kalau ke loket. Paling favorit di tempat merak putih 🐦💕	Netral	Negatif	Negatif
20	Gak cukup sehari kalo maen kesini karena luas banget dengan koleksi hewan yg banyak. Selain itu ada wahana bermain dan kolam renang anak. Wahana di dalam mayoritas sudah include di tiket masuk jadi gak perlu bayar lagi. Beberapa wahana ...	Positif	None	None
21	Tempat wisata yang cocok untuk keluarga. Saya mendatangi DINO PARK dan Millenial Glow Garden. Tempat yang sangat indah dan edukatif. Penerapan protokol kesehatan sangat ketat. Tersedia juga penyewaan kursi roda gratis untuk lansia. Memiliki konsep yang sangat millenial dan keren..	Positif	Positif	None
22	Alhamdulillah JP2 sudah dibuka kembali. Wahana2nya banyak yg baru, apa mungkin krn sy lama gak ke sini lagi ya jd gak tau. Protokol	Positif	Netral	None

No	review_teks	atraksi	amenitas	aksesibilitas
	kesehatan jg sangat diperhatikan mulai pintu masuk sampai di dalam, hingga keluar.			
23	Pas kesana rame full mengantri parah,panas dan tempat penuh sampai gk ada tempat buat duduk.semua wahana juga penuh antri ² berdesakan ,kolam anak penuh .. ada beberapa wahana yg rusak tidak bisa dinaiki,, trus juga banyak yg diperbaiki ...	Negatif	None	None
24	Tujuan perjalanan liburan ini menggunakan konsep edukatif yang menggabungkan dunia museum dinosaurus di abad-abad dan dunia legendaris alat musik dan seniman, negara dan pemimpin mereka, film dan seniman, olahraga dan seniman. Pengalaman liburan sangat menyenangkan. Buka dari 11 siang hingga 10 malam pada akhir pekan. Silakan menikmati dunia legendaris museum pertama, maka dunia museum dinosaurus. Ada juga berbagai kuliner, musik live dan merchandise. Hmmm bisa dalam satu hari tidak cukup untuk menikmati keseluruhan, jadi bagikan waktu Anda dengan baik. Nikmati bersama keluarga, kerabat dan teman-teman ...	Netral	Netral	Netral

Lampiran 3 : Perbandingan Anotasi LLM vs Pakar

No	Expert Person			LLM (Claude)		
	Atraksi	Aksesibilitas	Amenitas	Atraksi	Aksesibilitas	Amenitas
1	Positif	Positif	Positif	Positif	None	Positif
2	Positif	None	None	Positif	None	Netral
3	Positif	None	None	Positif	None	None
4	Positif	None	Positif	Positif	None	Positif
5	Positif	None	None	Positif	None	None
6	Netral	None	Positif	Positif	None	Positif

No	Expert Person			LLM (Claude)		
	Atraksi	Aksesibilitas	Amenitas	Atraksi	Aksesibilitas	Amenitas
7	Positif	None	Positif	Positif	None	Positif
8	Positif	None	Positif	Positif	Positif	Positif
9	Positif	Negatif	None	Positif	None	None
10	Positif	None	Positif	Positif	None	Positif
11	Positif	None	None	Positif	None	None
12	Positif	None	Positif	Netral	None	Positif
13	None	Positif	Positif	None	Positif	Positif
14	None	Netral	Positif	Positif	None	Positif
15	Positif	Positif	Positif	Positif	None	Positif
16	Positif	None	Netral	Positif	None	Positif
17	Positif	None	Positif	Positif	None	Positif
18	Positif	None	None	Positif	None	None
19	Netral	Negatif	Negatif	Netral	Negatif	Negatif
20	Positif	None	None	Positif	None	None
21	Positif	None	Positif	Positif	None	Positif
22	Positif	None	Netral	Positif	None	Netral
23	Negatif	None	None	Negatif	None	None
24	Netral	Netral	Netral	Positif	None	Netral

Lampiran 4 : Cuplikan Hasil Perbandingan

No	atraksi	aksesibilitas	amenitas
1	1	0	1
2	1	1	0
3	1	1	1
4	1	1	1
5	1	1	1
6	0	1	1

No	atraksi	aksesibilitas	amenitas
7	1	1	1
8	1	0	1
9	1	0	1
10	1	1	1
11	1	1	1
12	0	1	1
13	1	1	1
14	0	0	1
15	1	0	1
16	1	1	0
17	1	1	1
18	1	1	1
19	1	1	1
20	1	1	1
21	1	1	1
22	1	1	1
23	1	1	1
24	0	0	1



QR CODE DATASET
PENELITIAN