

Agung Kharisma Hidayah, S.Kom., M.Kom
Dandi Sunardi, M.Kom
Eka Sahputra, S.Kom., M.Kom



ANALISIS SENTIMEN

di Era Digital Konsep, Algoritma,
dan Studi Kasus
Media Sosial



ANALISIS SENTIMEN

di Era Digital Konsep, Algoritma,
dan Studi Kasus
Media Sosial

Agung Kharisma Hidayah, S.Kom., M.Kom
Dandi Sunardi, M.Kom
Eka Sahputra, S.Kom., M.Kom



**ANALISIS SENTIMEN DI ERA DIGITAL:
KONSEP, ALGORITMA, DAN STUDI KASUS MEDIA SOSIAL**

Ditulis oleh:

Agung Kharisma Hidayah, S.Kom., M.Kom.
Dandi Sunardi, M.Kom.
Eka Sahputra, S.Kom., M.Kom.

Diterbitkan, dicetak, dan didistribusikan oleh

PT Literasi Nusantara Abadi Grup

Perumahan Puncak Joyo Agung Residence Blok B11 Merjosari

Kecamatan Lowokwaru Kota Malang 65144

Telp : +6285887254603, +6285841411519

Email: literasinusantaraofficial@gmail.com

Web: www.penerbitlitnus.co.id

Anggota IKAPI No. 340/JTI/2022



Hak Cipta dilindungi oleh undang-undang. Dilarang mengutip
atau memperbanyak baik sebagian ataupun keseluruhan isi buku
dengan cara apa pun tanpa izin tertulis dari penerbit.

Cetakan I, Juli 2025

Perancang sampul: Bagus Aji Saputra

Penata letak: Noufal Fahriza

ISBN : 978-634-234-386-9

xii + 94 hlm., 15,5x23 cm.

©Juli 2025



KATA PENGANTAR

Di era digital seperti saat ini, opini publik tidak lagi hanya tercermin melalui survei atau forum resmi, melainkan juga melalui jutaan ekspresi spontan di media sosial. *Twitter, Facebook, Instagram*, dan *platform* daring lainnya telah menjadi cermin masyarakat digital yang memancarkan emosi, harapan, kritik, bahkan keresahan. Dalam arus informasi yang begitu deras, menjadi penting bagi peneliti, akademisi, dan pengambil kebijakan untuk memahami sentimen yang berkembang secara lebih cepat, sistematis, dan terukur.

Buku ini hadir sebagai respon atas kebutuhan tersebut. Dengan judul “*Analisis Sentimen di Era Digital: Konsep, Algoritma, dan Studi Kasus Media Sosial*”, kami berusaha menggabungkan landasan teoritis dengan pendekatan teknikal yang aplikatif. Pembaca tidak hanya akan diajak memahami konsep dasar seperti *sentiment analysis*, *natural language processing (NLP)*, dan *machine learning*, tetapi juga akan diperkenalkan dengan studi-studi kasus nyata yang mencerminkan dinamika sosial di dunia maya.

Kami memilih pemindahan Ibu Kota Negara (IKN) sebagai salah satu studi kasus utama dalam buku ini karena isu tersebut memicu diskusi yang luas dan beragam di masyarakat. Kami juga menyertakan berbagai contoh analisis sentimen lain di platform media sosial dan toko aplikasi, agar pembaca dapat melihat bagaimana teknik yang sama diterapkan pada konteks yang berbeda.

Buku ini ditujukan untuk mahasiswa, dosen, peneliti, praktisi data, dan siapa pun yang ingin memahami atau mengembangkan sistem analisis opini publik secara kuantitatif. Kami menyadari bahwa masih banyak ruang untuk penyempurnaan, baik dari sisi kedalaman

bahasan maupun keluasan aplikasinya. Namun, kami berharap buku ini dapat menjadi pijakan awal yang kokoh bagi pembaca untuk mengeksplorasi lebih lanjut potensi besar analisis sentimen di era digital ini.

Akhir kata, kami menyampaikan terima kasih kepada semua pihak yang telah mendukung penulisan buku ini. Semoga karya ini dapat memberi manfaat dan memperkaya khazanah literatur di bidang ilmu komputer, sosiologi digital, dan kebijakan publik berbasis data.

Bengkulu, Juli 2025

Penulis

Agung Kharisma Hidayah, S.Kom., M.Kom.

DAFTAR ISI

Kata Pengantar	iii
Daftar Isi.....	v
Daftar Gambar	ix
Daftar Tabel	xi

BAGIAN 1

EVOLUSI OPINI PUBLIK DI ERA DIGITAL..... 1

- A. Dari Forum Diskusi ke Platform Sosial 1
- B. Media Sosial sebagai Ruang Ekspresi Publik.....2
- C. Dinamika Waktu Nyata dan Kecepatan Opini3
- D. Tantangan dalam Menangkap Persepsi Publik.....4

BAGIAN 2

KOMPLEKSITAS BAHASA DI MEDIA SOSIAL..... 7

- A. Bahasa Tak Baku, *Slang*, dan *Emoticon*7
- B. Ambiguitas Makna dan *Code-Mixing*8
- C. Kendala Volume dan Arus Data Tinggi.....9
- D. Ketidakseimbangan Kategori Sentimen.....9
- E. Keterbatasan Model dan Representasi Teks 10

BAGIAN 3

MEMAHAMI ANALISIS SENTIMEN DAN PENDEKATANNYA.....13

- A. Pendekatan *Lexicon-Based* 13

B. Pendekatan <i>Machine Learning-Based</i>	14
C. Pendekatan <i>Hybrid</i>	15
D. Aplikasi Analisis Sentimen di Dunia Nyata	15

BAGIAN 4

DASAR PEMROSESAN BAHASA ALAMI (*NATURAL LANGUAGE PROCESSING*)17

A. Apa Itu Pemrosesan Bahasa Alami?	17
B. <i>Preprocessing</i> Teks.....	18
C. Tantangan Spesifik Bahasa Indonesia.....	19

BAGIAN 5

REPRESENTASI DATA TEKS DAN TEKNIK FITURISASI21

A. Mengubah Teks Menjadi Data Numerik	21
B. <i>Bag-of-Words</i> (BoW).....	23
C. Term Frequency-Inverse Document Frequency (TF-IDF).....	25
D. <i>Word Embeddings</i> : Memahami Makna Kata secara Semantik.....	28
E. BERT dan Model <i>Transformer</i>	30

BAGIAN 6

PENYEIMBANGAN DATA DAN ALGORITMA KLASIFIKASI..... 35

A. Permasalahan Ketimpangan Kategori.....	35
B. Teknik Penyeimbangan: <i>SMOTE</i> dan Alternatifnya	37
C. Alternatif <i>SMOTE</i> : Strategi Lain dalam Penyeimbangan Data	39
D. Pengenalan Algoritma Klasifikasi	40
E. Evaluasi Performa Model	43

BAGIAN 7

PENERAPAN ANALISIS SENTIMEN PADA ISU SOSIAL.....	47
A. Topik dan Konteks: Pemindahan Ibu Kota Negara.....	48
B. Trategi Pengambilan Data dari Twitter	50
C. Proses Pelabelan Sentimen.....	51
D. Distribusi Sentimen dalam Dataset	53

BAGIAN 8

IMPLEMENTASI ANALISIS SENTIMEN PADA STUDI KASUS	55
A. Pra-pemrosesan Data Teks.....	56
B. Pembentukan Fitur dengan TF-IDF.....	57
C. Penyeimbangan Data Menggunakan SMOTE.....	59
D. Pelatihan Model Klasifikasi dengan <i>Random Forest</i>	61
E. Evaluasi Kinerja Model.....	63
F. Visualisasi dan Interpretasi Hasil	65
G. Hasil Evaluasi Dan Interpretasi	67
H. Perbandingan dengan Pendekatan Lain	69
I. Analisis Kata Kunci Berdasarkan Sentimen	70

BAGIAN 9

ILUSTRASI DALAM PENERAPAN SENTIMEN DIGITAL	73
A. Sentimen Pengguna Aplikasi Vidio	73
B. Ulasan Pengguna Blibli dan Teknik Gabungan	75
C. Perbandingan Persepsi Pengguna Shopee vs TikTok Shop.....	78
D. Analisis Sentimen dalam Isu Sosial Viral	79

BAGIAN 10

PANDUAN PRAKTIS IMPLEMENTASI SISTEM ANALISIS SENTIMEN	83
A. Pendahuluan.....	83
B. Contoh kode program.....	84
C. Panduan Eksekusi Program	86
D. Penutup.....	86
Daftar Pustaka.....	89
Tentang Penulis	93



DAFTAR GAMBAR

Gambar 1.	Distribusi Sentimen Awal sebelum SMOTE	54
Gambar 2.	Perbandingan Distribusi Sentimen Sebelum dan Sesudah SMOTE.....	60
Gambar 3.	Diagram Alur Proses Pelatihan Model.....	61
Gambar 4.	Confusion Matrix – Random Forest	65
Gambar 5.	Word Cloud Sentimen Positif.....	68
Gambar 6.	Word Cloud Sentimen Negatif	68
Gambar 7.	Word Cloud Sentimen Positif Aplikasi Video.....	74
Gambar 8.	Word Cloud Sentimen Negatif Aplikasi Video	75
Gambar 9.	Word Cloud Ulasan Positif Blibli	77
Gambar 10.	Word Cloud Ulasan Negatif Blibli	77
Gambar 11.	Confusion matrix Naïve Bayes dan SVM	81

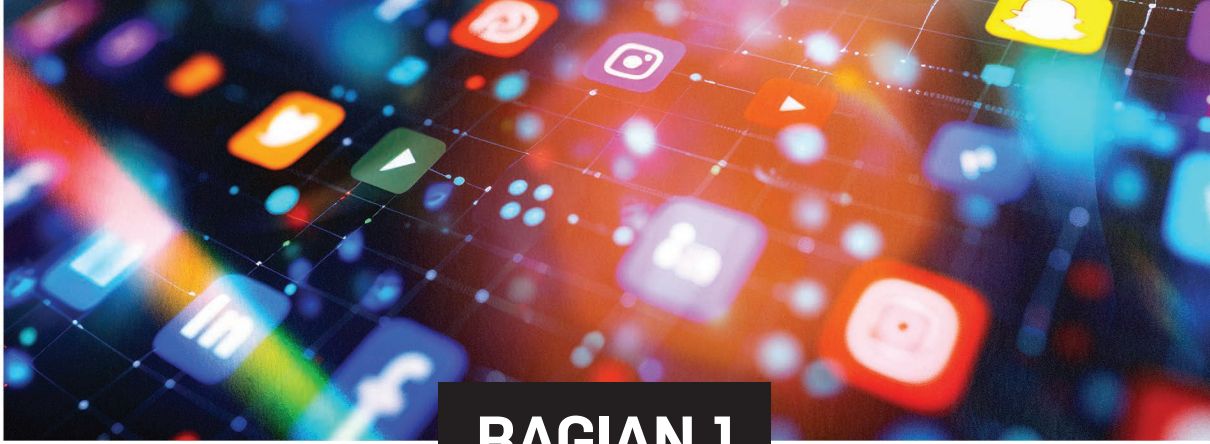


DAFTAR TABEL

Tabel 1. Contoh Matrik Frekuensi Kata..... 24

Tabel 2. Evaluasi Model Random Forest..... 64

Tabel 3. Kata Kunci Berdasarkan Aspek dan Sentimen..... 71



BAGIAN 1

EVOLUSI OPINI PUBLIK DI ERA DIGITAL

A. Dari Forum Diskusi ke Platform Sosial

Opini publik merupakan kekuatan sosial yang telah lama memainkan peran sentral dalam membentuk arah kebijakan, budaya, hingga sistem nilai suatu masyarakat. Dalam konteks klasik, opini publik berkembang melalui forum-forum diskusi formal maupun informal, seperti dewan rakyat di Athena Kuno, balai kota di masa feodal, hingga kolom opini di surat kabar. Pola interaksi bersifat hierarkis, terbatas oleh ruang dan waktu, serta cenderung dimonopoli oleh elite atau kelompok terdidik.

Namun, seiring masuknya teknologi digital ke ruang komunikasi publik, struktur ini mengalami pergeseran radikal. Internet menjadi pemantik utama transformasi, yang kemudian disusul oleh kehadiran media sosial. Dimulai dari *mailing list* dan *web forum* seperti *Usenet*, *Kaskus*, atau *Kompasiana*, publik mulai memiliki ruang alternatif untuk menyuarakan pendapat secara langsung. Tidak perlu lagi menunggu ruang redaksi media massa — siapa pun bisa menjadi *produsen opini*.

Memasuki era media sosial, perubahan ini menjadi semakin masif. Platform seperti Facebook, Twitter, YouTube, dan Instagram menghadirkan komunikasi dua arah yang tak mengenal batas geografis maupun otoritas tunggal. Media sosial membuka keran partisipasi tak terbatas: semua orang dapat membagikan ide, membentuk opini, memengaruhi wacana, bahkan memobilisasi tindakan kolektif.

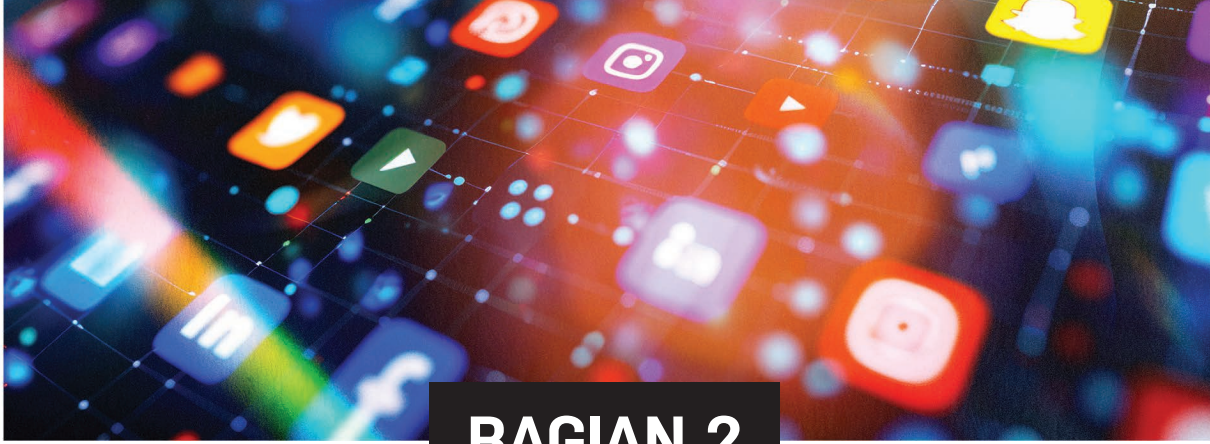
Yang menarik, media sosial tidak hanya merepresentasikan pendapat, tetapi juga merekam emosi. Ketika seseorang menuliskan kekecewaannya terhadap layanan publik, menyuarakan dukungan terhadap gerakan sosial, atau sekadar membagikan kegelisahan pribadi atas isu nasional, yang terekam bukan hanya informasi verbal, tetapi juga nuansa afektif yang melatarbelakanginya. Dalam hal ini, media sosial telah bertransformasi menjadi *ruang psiko-sosial digital* — tempat di mana identitas, opini, dan emosi berinteraksi secara simultan.

Hal ini menghadirkan tantangan sekaligus peluang. Tantangan karena volume dan keberagaman opini yang muncul sangat besar dan dinamis. Peluang karena data tersebut dapat dianalisis untuk memahami sentimen kolektif masyarakat secara *real time*, memberikan masukan bagi para pengambil kebijakan, pengusaha, dan peneliti sosial.

B. Media Sosial sebagai Ruang Ekspresi Publik

Media sosial berperan bukan hanya sebagai sarana komunikasi, tetapi juga sebagai arena diskursif di mana opini masyarakat dibentuk, diperdebatkan, dan disebarluaskan. Di sinilah *netizen* membentuk persepsi kolektif tentang suatu isu: melalui *retweet*, komentar, *hashtag*, dan bahkan meme.

Sifat terbuka dan partisipatif media sosial menciptakan ekosistem opini yang inklusif sekaligus tak terkontrol. Berbeda dengan media konvensional yang umumnya melalui proses penyuntingan, opini



BAGIAN 2

KOMPLEKSITAS BAHASA DI MEDIA SOSIAL

A. Bahasa Tak Baku, *Slang*, dan *Emoticon*

Bahasa yang digunakan dalam media sosial sangat berbeda dari bahasa dalam dokumen formal atau komunikasi resmi. Media sosial menumbuhkan jenis komunikasi yang lincah, informal, penuh improvisasi, dan sangat kontekstual. Pengguna sering menggunakan bahasa sehari-hari, singkatan, hingga *slang* yang unik bagi komunitas tertentu. Misalnya, dalam bahasa Indonesia, ungkapan seperti “mager”, “gabut”, “auto kaya”, atau “gaskeun” adalah bagian dari dinamika bahasa digital yang tidak ditemukan dalam kamus konvensional.

Selain kata-kata tidak baku, pengguna juga sering menggunakan *emoticon* dan emoji untuk mengekspresikan perasaan mereka. Misalnya, emoji 🤔, 😭, atau 😂 memiliki makna sentimen yang kuat, meski tidak berbentuk teks. Bahkan dalam beberapa kasus, sebuah emoji dapat menjadi inti dari pesan itu sendiri, menggantikan seluruh kalimat atau menambahkan konteks emosional yang tidak bisa dibaca hanya dari teks polos.

Ketidakkakuan ini menimbulkan tantangan besar dalam analisis sentimen otomatis. Sistem komputer sering kali dirancang untuk

memproses bahasa baku dan standar. Ketika dihadapkan dengan teks informal, sistem tersebut bisa salah mengartikan konteks atau gagal mengenali makna sebenarnya dari pesan. Oleh karena itu, diperlukan strategi pemrosesan awal (*preprocessing*) yang mampu menormalisasi atau menginterpretasikan bentuk-bentuk bahasa informal ini agar bisa diproses secara valid oleh sistem.

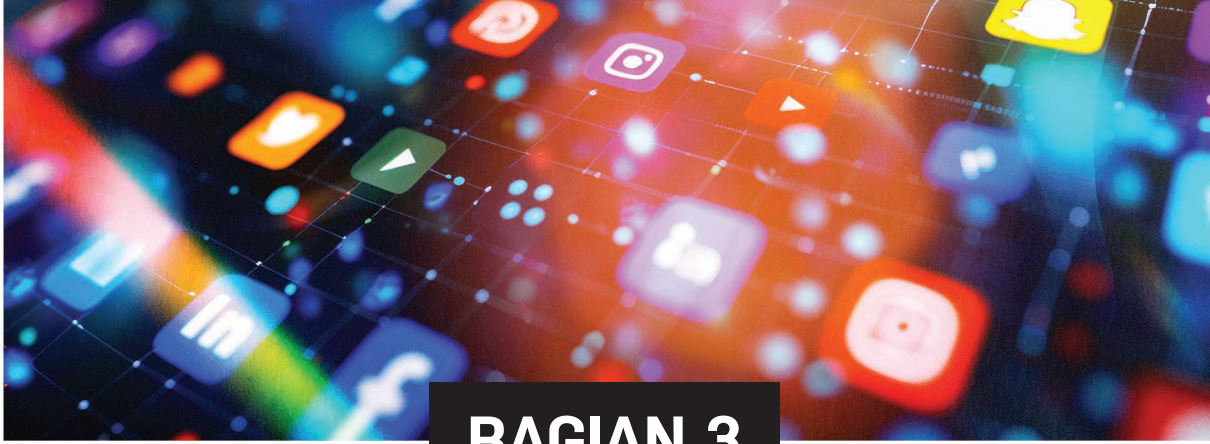
B. Ambiguitas Makna dan *Code-Mixing*

Salah satu ciri khas komunikasi di media sosial adalah munculnya banyak ambiguitas — baik dalam makna kata maupun dalam nada penyampaian. Pengguna sering menggunakan bahasa sarkastik, ironi, atau humor yang sulit dipahami jika hanya dilihat dari teks literal. Contohnya, kalimat “Wah hebat banget pelayanan kamu, mantap!” bisa menjadi pujian tulus atau justru sindiran tajam, tergantung konteks percakapan atau emosi penulisnya.

Tak jarang pula, pengguna mencampur dua atau lebih bahasa dalam satu unggahan atau komentar — sebuah fenomena yang dikenal sebagai *code-mixing* atau *code-switching*. Di Indonesia, percampuran antara Bahasa Indonesia dan Bahasa Inggris sangat umum ditemukan, seperti dalam kalimat “Dia tuh literally ghosting banget, fix toxic.” Kalimat semacam ini menyulitkan sistem *Natural Language Processing (NLP)* yang tidak dilatih untuk memproses teks multibahasa secara bersamaan.

Selain itu, pengguna juga sering menyingkat kata (misalnya “gk”, “sy”, “bgt”) atau menggunakan gaya penulisan kreatif seperti huruf kapital berlebihan, pengulangan huruf (“baguuuuuusss”), atau simbol tambahan (“-_-”, “!!!”, dll) yang membawa muatan emosional tertentu.

Mendeteksi makna sesungguhnya dari ungkapan-ungkapan tersebut memerlukan sistem analisis yang tidak hanya mengenali kata, tetapi juga mampu menangkap konteks, nada, serta nuansa budaya dari penutur bahasa tersebut.



BAGIAN 3

MEMAHAMI ANALISIS SENTIMEN DAN PENDEKATANNYA

Secara umum, terdapat dua pendekatan utama dalam melaksanakan analisis sentimen: pendekatan berbasis *lexicon* dan pendekatan berbasis *machine learning*.

A. Pendekatan *Lexicon-Based*

Pendekatan ini menggunakan kamus atau daftar kata yang telah diberi label sentimen, baik positif maupun negatif. Setiap kata dalam teks akan dibandingkan dengan daftar tersebut, dan hasilnya digunakan untuk menentukan keseluruhan sentimen teks.

Contoh sederhana: jika dalam satu kalimat terdapat lebih banyak kata dari daftar positif (“bagus”, “mantap”, “cepat”) daripada kata dari daftar negatif (“buruk”, “lambat”, “mengecewakan”), maka kalimat tersebut diasumsikan memiliki sentimen positif.

Kelebihan:

1. Tidak memerlukan data pelatihan.
2. Mudah diimplementasikan dan cepat diproses.

Kekurangan:

1. Sensitif terhadap konteks; gagal mengenali ironi atau sarkasme.
2. Terbatas jika tidak ada daftar kata yang relevan dengan domain tertentu.

B. Pendekatan *Machine Learning-Based*

Pendekatan ini menggunakan algoritma *machine learning* untuk mempelajari pola sentimen dari data berlabel. Sistem dilatih dengan sejumlah besar data teks yang telah diberi anotasi sentimen, lalu digunakan untuk memprediksi sentimen dari data baru.

Tahapannya mencakup:

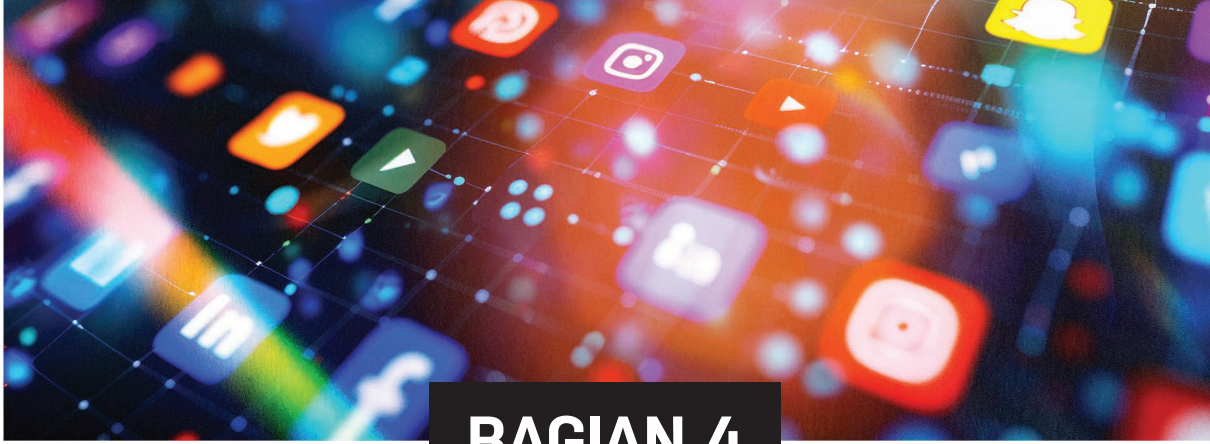
1. *Preprocessing* teks: seperti *case folding*, penghapusan simbol, *tokenisasi*, dan *stemming*.
2. Ekstraksi fitur: biasanya menggunakan metode seperti *TF-IDF* atau *word embeddings*.
3. Pelatihan model klasifikasi: dengan algoritma seperti *Naïve Bayes*, *Support Vector Machine (SVM)*, atau *Random Forest*.
4. Evaluasi performa model dengan metrik seperti *accuracy*, *precision*, *recall*, dan *F1-score*.

Kelebihan:

1. Lebih adaptif dan akurat untuk berbagai jenis teks dan domain.
2. Dapat menangkap pola kompleks yang tidak terlihat secara eksplisit.

Kekurangan:

1. Membutuhkan data latih yang cukup dan berkualitas.
2. Bisa menjadi *overfitting* jika tidak disusun dengan benar.



BAGIAN 4

DASAR PEMROSESAN BAHASA ALAMI (*NATURAL LANGUAGE PROCESSING*)

A. Apa Itu Pemrosesan Bahasa Alami?

Natural Language Processing (NLP) adalah cabang dari ilmu komputer yang berfokus pada interaksi antara komputer dan bahasa manusia. Tujuannya adalah memungkinkan mesin untuk “memahami”, menganalisis, dan menghasilkan bahasa yang digunakan oleh manusia secara alami. Dalam konteks analisis sentimen, *NLP* berperan sebagai fondasi utama dalam mengekstraksi makna dari teks.

Berbeda dengan bahasa formal seperti bahasa pemrograman, bahasa alami memiliki karakteristik yang kompleks: ambigu, kontekstual, penuh nuansa budaya, dan sering kali tidak mengikuti kaidah gramatikal yang ketat. Tantangan inilah yang membuat *NLP* menjadi bidang yang sangat dinamis dan terus berkembang.

Penerapan *NLP* sangat luas: dari mesin pencari (*search engine*), penerjemah otomatis (*machine translation*), chatbot, hingga sistem rekomendasi. Untuk keperluan analisis sentimen, *NLP* menyediakan teknik untuk menyiapkan teks mentah agar bisa dibaca oleh sistem — sebuah proses yang dikenal dengan istilah *text preprocessing*.

B. *Preprocessing* Teks

Dalam analisis sentimen berbasis teks, *preprocessing* merupakan tahap awal yang sangat penting untuk meningkatkan kualitas data sebelum digunakan dalam pemodelan. Teks yang berasal dari media sosial seperti Twitter cenderung tidak terstruktur, mengandung banyak kata tidak baku, singkatan, simbol, emoji, tautan, dan elemen informal lainnya. Oleh karena itu, *preprocessing* bertujuan untuk membersihkan dan menormalkan teks agar dapat diproses secara komputasional.

Tahapan *preprocessing* teks dalam penelitian ini mencakup beberapa langkah berikut:

1. *Case Folding*

Case folding adalah proses mengubah seluruh huruf dalam teks menjadi huruf kecil (*lowercase*). Misalnya, kata “Indonesia”, “INDONESIA”, dan “indonesia” dianggap berbeda oleh komputer, sehingga perlu diseragamkan menjadi “indonesia”.

2. Penghapusan Simbol dan Karakter Khusus

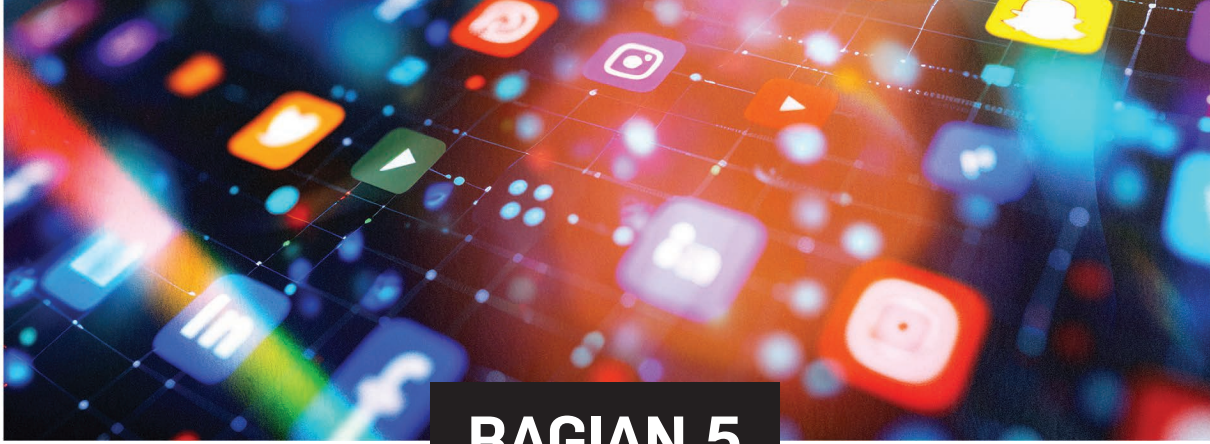
Simbol seperti tanda baca, angka, emoji, tagar (#), mention (@), dan tautan (<https://...>) biasanya tidak memiliki nilai informatif dalam analisis sentimen dan harus dihapus, kecuali jika dibutuhkan untuk fitur khusus.

3. Tokenisasi

Tokenisasi adalah proses memecah teks menjadi unit-unit kata (token). Misalnya, kalimat “IKN terlalu mahal” diubah menjadi daftar token: [‘ikn’, ‘terlalu’, ‘mahal’]. Ini memudahkan proses analisis dan ekstraksi fitur.

4. *Stopword Removal*

Stopword adalah kata-kata umum yang tidak memiliki makna penting dalam analisis, seperti “yang”, “dan”, “di”, “ke”, dan sebagainya. Kata-kata ini dihapus agar fitur yang digunakan lebih relevan dan tidak mendominasi hasil klasifikasi.



BAGIAN 5

REPRESENTASI DATA TEKS DAN TEKNIK FITURISASI

Setelah data teks melalui tahap *preprocessing*, langkah selanjutnya dalam analisis sentimen adalah mengubah data teks menjadi representasi numerik yang dapat diproses oleh algoritma pembelajaran mesin. Proses ini dikenal sebagai ekstraksi fitur (*feature extraction*). Dalam konteks pemrosesan bahasa alami, fitur yang umum digunakan berasal dari kata-kata atau frasa yang muncul dalam teks. Beberapa pendekatan paling populer untuk ekstraksi fitur adalah *Bag-of-Words* (BoW), *Term Frequency-Inverse Document Frequency* (TF-IDF), *Word Embeddings* dan

A. Mengubah Teks Menjadi Data Numerik

Bahasa manusia adalah sesuatu yang kaya, kompleks, dan ambigu. Dalam bentuk aslinya — kalimat, paragraf, atau dokumen — teks tidak dapat diproses secara langsung oleh komputer. Komputer tidak memahami kata “bagus” atau “mengecewakan” seperti manusia memahami maknanya. Untuk bisa melakukan analisis sentimen secara otomatis, langkah pertama yang sangat krusial adalah **mengubah teks menjadi representasi numerik** yang dapat diproses oleh algoritma pembelajaran mesin (*machine learning*). Proses ini

dikenal sebagai *feature extraction*, yaitu proses mengubah kata atau dokumen menjadi fitur-fitur numerik yang mewakili informasi penting dalam teks.

Transformasi ini bukan sekadar proses teknis, tetapi merupakan inti dari seluruh sistem analisis sentimen. Jika representasi teks tidak akurat atau gagal menangkap makna kata dan relasinya, maka model analisis yang dibangun di atasnya juga akan memberikan hasil yang tidak memuaskan. Oleh karena itu, memahami bagaimana teks direpresentasikan secara numerik menjadi dasar yang sangat penting dalam membangun sistem klasifikasi sentimen yang handal.

Mengapa Teks Harus Diubah Menjadi Angka?

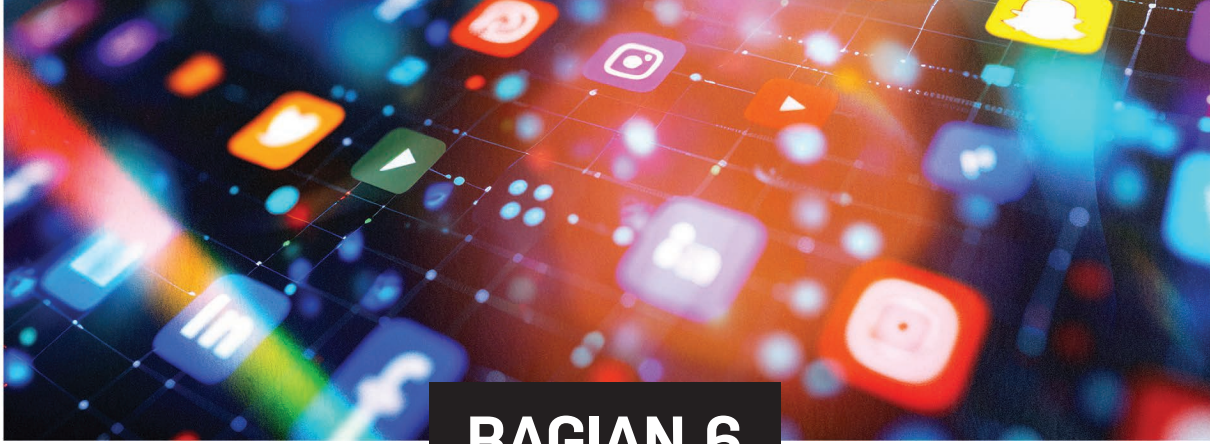
Komputer dan algoritma *machine learning* hanya dapat bekerja dengan angka. Ketika kita berbicara tentang klasifikasi teks — misalnya membedakan ulasan positif dan negatif — maka kita membutuhkan *fitur* (ciri-ciri) dalam bentuk angka yang dapat dimasukkan ke dalam model klasifikasi seperti *Naïve Bayes*, *Support Vector Machine*, atau *Random Forest*. Fitur-fitur tersebut harus dapat mencerminkan aspek penting dari sebuah dokumen: misalnya kata apa yang muncul, seberapa sering kata tersebut digunakan, bagaimana konteks penggunaannya, dan sebagainya.

Tanpa representasi numerik yang baik, model tidak akan bisa belajar membedakan antara kalimat “saya puas dengan pelayanannya” dan “saya sangat kecewa dengan pelayanan buruknya.”

Setiap pendekatan memiliki kekuatan dan keterbatasannya sendiri. Pendekatan yang dipilih harus disesuaikan dengan jenis data, kompleksitas masalah, dan sumber daya yang tersedia.

Meskipun proses transformasi teks ke dalam bentuk numerik telah berkembang sangat pesat, ada sejumlah tantangan yang tetap perlu diperhatikan:

1. **Ukuran Dimensi:** Representasi berbasis kata seperti BoW atau TF-IDF bisa menghasilkan vektor dengan ribuan dimensi, yang dapat memperlambat proses pelatihan model.



BAGIAN 6

PENYEIMBANGAN DATA DAN ALGORITMA KLASIFIKASI

A. Permasalahan Ketimpangan Kategori

Salah satu tantangan umum yang sering muncul dalam analisis sentimen — terutama saat mengambil data dari media sosial atau ulasan daring — adalah **ketimpangan kategori** (*class imbalance*). Ketimpangan ini terjadi ketika jumlah data pada satu kelas sentimen jauh lebih besar dibandingkan kelas lainnya. Misalnya, dalam suatu dataset ulasan aplikasi, terdapat ribuan ulasan positif, namun hanya puluhan ulasan negatif. Atau dalam diskusi tentang kebijakan publik, sebagian besar komentar mungkin bersifat netral, sementara opini yang kritis atau negatif jumlahnya sangat sedikit.

Masalah ini tampak sederhana, namun bisa berdampak besar pada kinerja model *machine learning* yang dilatih menggunakan data tersebut. Secara statistik, model cenderung “bermain aman” dengan selalu memprediksi kelas mayoritas. Ini bisa menghasilkan **akurasi tinggi secara keseluruhan**, tetapi **kegagalan total dalam mengenali kelas minoritas**, yang justru sering kali paling penting untuk dianalisis.

Sebagai contoh, dalam konteks pemantauan opini publik terhadap isu pemindahan Ibu Kota Negara (IKN), ulasan atau komentar yang bersifat negatif mungkin menjadi indikator awal dari resistensi publik, potensi konflik, atau ketidakpuasan kebijakan. Jika model gagal mendeteksi opini-opini minoritas ini karena jumlahnya yang terlalu kecil, maka sistem analisis tidak akan mampu memberikan gambaran yang akurat dan bermanfaat bagi pengambil kebijakan atau pihak-pihak yang berkepentingan.

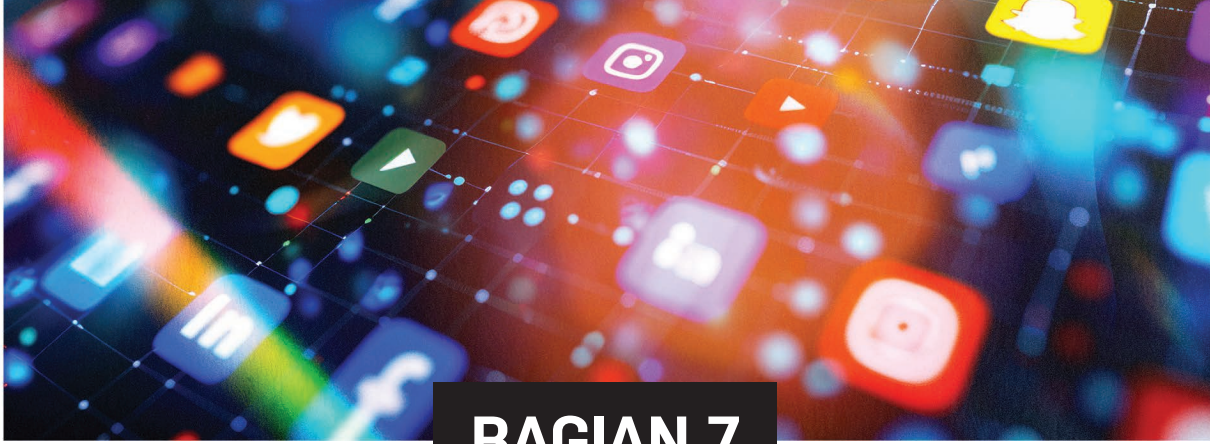
Model klasifikasi bekerja berdasarkan pola yang dilatih dari data. Jika sebagian besar data yang digunakan untuk pelatihan berasal dari satu kategori (misalnya: *positif*), maka model akan belajar untuk mengenali pola tersebut dengan sangat baik, tetapi gagal dalam mengenali pola dari kategori lain (misalnya: *negatif* atau *netral*). Akibatnya:

1. Prediksi akan bias ke kelas mayoritas
2. *Recall* dan *precision* untuk kelas minoritas menjadi sangat rendah
3. Evaluasi kinerja model menjadi menyesatkan bila hanya melihat akurasi keseluruhan

Situasi ini tidak hanya mengganggu validitas sistem, tetapi juga membahayakan pengambilan keputusan berbasis data.

Bayangkan sebuah model yang dilatih untuk mengklasifikasi sentimen pengguna terhadap aplikasi tertentu. Dari 10.000 data, sebanyak 9.500 adalah ulasan positif, dan hanya 500 negatif. Model yang *selalu memprediksi positif* akan memiliki akurasi 95%, tetapi sebenarnya **tidak pernah berhasil menangkap sentimen negatif**, yang sering kali justru lebih relevan untuk evaluasi produk.

Ketimpangan kategori seperti ini bukan hanya masalah teknis, tetapi juga masalah strategis. Dalam pengambilan keputusan berbasis data, suara minoritas sering kali membawa nilai yang sangat penting. Oleh karena itu, mengatasi ketimpangan data menjadi keharusan dalam setiap sistem analisis sentimen yang serius.



BAGIAN 7

PENERAPAN ANALISIS SENTIMEN PADA ISU SOSIAL

Pemindahan Ibu Kota Negara (IKN) dari Jakarta ke Kalimantan Timur merupakan kebijakan strategis nasional yang mencerminkan ambisi besar pemerintah dalam mewujudkan pemerataan pembangunan dan desentralisasi kekuasaan. Keputusan ini menuai respons luas dari masyarakat, baik dalam bentuk dukungan maupun penolakan, yang terekam dengan jelas melalui berbagai platform media sosial. Salah satu kanal utama yang mencerminkan reaksi publik secara real-time dan terbuka adalah **Twitter**, di mana masyarakat dari berbagai latar belakang dapat menyampaikan opini, kritik, atau harapan mereka terhadap kebijakan tersebut.

Kebijakan pemindahan IKN bukan sekadar proyek infrastruktur. Ia membawa serta dimensi politik, ekonomi, sosial, bahkan ideologis. Isu-isu seperti pembiayaan proyek yang besar, urgensi pemindahan di tengah tantangan lain seperti pandemi dan inflasi, hingga kekhawatiran terhadap dampak lingkungan menjadi bagian dari percakapan digital yang berlangsung sejak wacana ini diumumkan. Di sisi lain, argumen positif juga muncul, seperti peluang pertumbuhan kawasan timur Indonesia, pemerataan ekonomi, dan solusi jangka panjang atas kepadatan Jakarta. Kompleksitas opini

tersebut menjadikan pemindahan IKN sebagai topik yang ideal untuk dianalisis menggunakan pendekatan analisis sentimen berbasis *machine learning*.

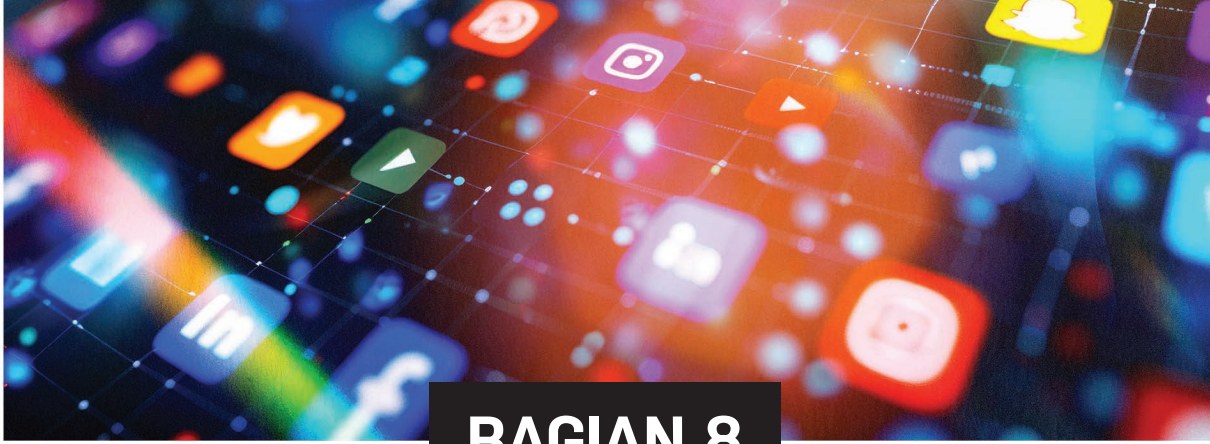
Studi kasus ini mengambil data dari Twitter sebagai sumber utama opini publik. Dengan membatasi cakupan pada kata kunci relevan seperti “IKN”, “ibu kota negara”, “pindah ibu kota”, dan “Nusantara”, penelitian ini mengumpulkan ribuan *tweet* yang berisi pendapat, respons, dan diskusi warganet terkait kebijakan tersebut. Data ini selanjutnya diolah dan dianalisis untuk mengungkap pola sentimen masyarakat, distribusi opini, serta respons emosional terhadap berbagai aspek dari proyek pemindahan IKN.

Bab ini akan menjelaskan secara terperinci konteks pemindahan IKN, alasan pemilihan topik sebagai studi kasus, serta bagaimana data dikumpulkan dari Twitter untuk dianalisis. Pemahaman mendalam terhadap isu ini menjadi penting agar analisis yang dilakukan tidak hanya bersifat teknis, tetapi juga kontekstual dan relevan secara sosial-politik. Dengan demikian, studi kasus ini berfungsi sebagai jembatan antara metodologi ilmiah dan realitas sosial yang dianalisis.

A. Topik dan Konteks: Pemindahan Ibu Kota Negara

Pemindahan Ibu Kota Negara dari Jakarta ke Kalimantan Timur merupakan proyek ambisius pemerintah Indonesia yang menuai sorotan luas dari berbagai pihak. Rencana ini secara resmi diumumkan oleh Presiden Joko Widodo pada tahun 2019, dengan latar belakang kebutuhan mendesak untuk mengurangi beban Jakarta sebagai pusat pemerintahan, ekonomi, dan populasi. Kota baru yang diberi nama “**Nusantara**” diharapkan dapat menjadi simbol pemerataan pembangunan nasional dan mendorong pertumbuhan kawasan timur Indonesia.

Meski dilandasi oleh visi jangka panjang, kebijakan pemindahan ini juga memicu beragam reaksi, terutama di ruang publik digital.



BAGIAN 8

IMPLEMENTASI ANALISIS SENTIMEN PADA STUDI KASUS

Bagian ini membahas secara rinci langkah-langkah implementasi sistem analisis sentimen terhadap opini publik mengenai pemindahan Ibu Kota Negara (IKN) yang bersumber dari data media sosial Twitter. Proses implementasi dimulai dari tahap pra-pemrosesan data teks, dilanjutkan dengan transformasi data menjadi vektor numerik, penyeimbangan kelas sentimen menggunakan SMOTE, dan diakhiri dengan pembangunan serta evaluasi model klasifikasi menggunakan algoritma *Random Forest*.

Setiap tahapan dalam *pipeline* ini memiliki peran penting yang saling terkait. *Preprocessing* yang baik akan meningkatkan akurasi model karena menghilangkan *noise* dan membuat data lebih seragam. Transformasi data dengan teknik TF-IDF memungkinkan teks dianalisis dalam bentuk numerik, sedangkan SMOTE membantu menyeimbangkan representasi setiap kelas agar model tidak bias terhadap kelas mayoritas. Algoritma *Random Forest* kemudian digunakan sebagai alat klasifikasi utama yang diharapkan mampu memberikan hasil prediksi yang akurat dan stabil, bahkan dalam kondisi data yang kompleks dan tidak terstruktur seperti teks Twitter.

Penerapan metode-metode ini tidak hanya menjawab persoalan teknis dalam pengolahan data teks, tetapi juga mencerminkan strategi konseptual untuk menghasilkan model yang tidak hanya akurat secara statistik, tetapi juga adil dan representatif terhadap keberagaman opini masyarakat.

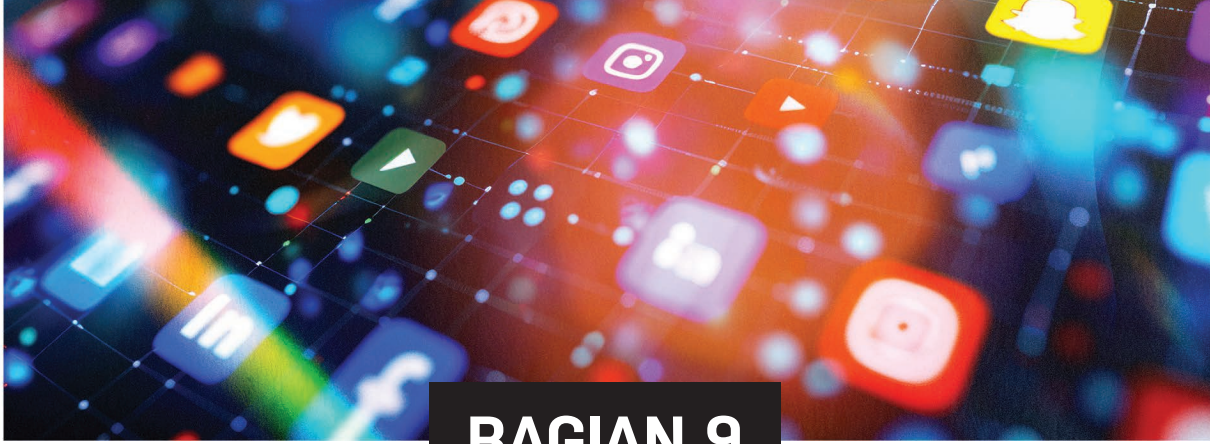
A. Pra-pemrosesan Data Teks

Tahap pertama dalam implementasi sistem analisis sentimen adalah pra-pemrosesan data teks (*text preprocessing*). Langkah ini sangat krusial karena data mentah dari media sosial, khususnya Twitter, umumnya mengandung banyak elemen yang tidak relevan atau mengganggu, seperti tagar, mention, emoji, URL, serta struktur bahasa yang informal. Jika tidak ditangani dengan baik, elemen-elemen ini dapat menyebabkan kekeliruan pada proses pembelajaran mesin, menurunkan akurasi model, dan menciptakan hasil yang bias atau tidak representatif.

Dalam penelitian ini, pra-pemrosesan dilakukan secara sistematis dengan beberapa tahap utama, yang seluruhnya diotomatisasi menggunakan skrip *Python*. Setiap tahap memiliki peran tersendiri dalam menyederhanakan, menyaring, dan menstandarkan teks agar siap digunakan sebagai input dalam tahap vektorisasi dan klasifikasi.

Langkah pertama adalah konversi huruf menjadi huruf kecil (*case folding*), untuk memastikan bahwa variasi kapitalisasi seperti “IKN” dan “ikn” tidak dianggap sebagai entitas berbeda. Selanjutnya dilakukan penghapusan karakter khusus, seperti simbol tanda baca, angka, dan emoji, serta penghilangan tautan (URL) dan mention (@username) yang tidak memberikan kontribusi bermakna terhadap sentimen.

Setelah itu, teks dipecah menjadi kata-kata individu melalui proses tokenisasi, sehingga setiap tweet diubah menjadi serangkaian token. Token-token ini kemudian diperiksa untuk menghapus *stopword*, yaitu kata-kata umum dalam bahasa Indonesia seperti



BAGIAN 9

ILUSTRASI DALAM PENERAPAN SENTIMEN DIGITAL

A. Sentimen Pengguna Aplikasi Vidio

Aplikasi *Vidio*, sebagai salah satu *platform streaming* digital populer di Indonesia, menghadirkan berbagai layanan siaran langsung olahraga, serial drama, film, hingga konten lokal premium. Seiring pertumbuhan pengguna aktifnya, *platform* ini juga mendapat banyak tanggapan publik dalam bentuk ulasan yang diposting melalui *Google Play Store*. Respons tersebut memuat persepsi, pengalaman, dan harapan pengguna terhadap performa dan kualitas layanan. Oleh karena itu, analisis sentimen terhadap ulasan ini menjadi penting untuk memahami pola kepuasan dan ketidakpuasan konsumen.

Penelitian ini mengadopsi pendekatan klasifikasi sentimen berbasis ***Support Vector Machine (SVM)***. Data dikumpulkan melalui teknik *web scraping* terhadap 2.000 ulasan pengguna aplikasi Vidio yang ditulis dalam rentang waktu 15 September hingga 21 November 2024. Label sentimen dibentuk secara otomatis berdasarkan sistem rating bintang: ulasan dengan rating 4 dan 5 dikategorikan sebagai positif, sementara rating 1–3 dianggap sebagai negatif.



BAGIAN 10

PANDUAN PRAKTIS IMPLEMENTASI SISTEM ANALISIS SENTIMEN

A. Pendahuluan

Bagian ini menyajikan *listing* program lengkap untuk melakukan analisis sentimen terhadap opini publik mengenai pemindahan Ibu Kota Negara (IKN), yang ditulis dalam bahasa pemrograman *Python*. Model yang digunakan adalah ***Random Forest***, dengan pendekatan vektorisasi teks menggunakan **TF-IDF** dan penyeimbangan kelas menggunakan **SMOTE**. Program ini dapat dijalankan di *Google Colab* atau *Jupyter Notebook* dan dirancang untuk dapat direplikasi secara mudah.

Setiap blok kode disertai dengan penjelasan singkat mengenai fungsinya. Di akhir setiap tahap, hasil olahan juga diekspor ke file CSV agar dapat digunakan untuk evaluasi, dokumentasi, maupun visualisasi lanjutan.

B. Contoh kode program

```
# === 1. Import Library ===
import pandas as pd
import numpy as np
import re
from sklearn.model_selection import train_
test_split
from sklearn.feature_extraction.text import
TfidfVectorizer
from sklearn.ensemble import
RandomForestClassifier
from sklearn.metrics import classification_
report, confusion_matrix
from imblearn.over_sampling import SMOTE

# === 2. Load Dataset ===
df = pd.read_csv("/content/Sentimen_
Pengguna_Twitter_Pada_Topik_IKN.csv")
df = df[['Sentimen', 'Tweet']].dropna()

# === 3. Preprocessing Teks ===
def clean_text(text):
    text = str(text).lower()
    text = re.sub(r"http\S+|@\w+|#\w+|[^a-
z\s]", "", text)
    text = re.sub(r"\s+", " ", text).strip()
    return text
df['Clean_Tweet'] = df['Tweet'].
apply(clean_text)
df.to_csv("01_Preprocessed_Data.csv",
index=False)
```

DAFTAR PUSTAKA

- A. Adi Bhirawa and U. Pradema Sanjaya, "From Data Imbalance to Precision: SMOTE-Driven Machine Learning for Early Detection of Kidney Disease," *INOVTEK Polbeng–Seri Inform.*, vol. 10, no. 1, pp. 514–525, 2025, doi: 10.35314/7jgimg64.
- A. F. Anjani, D. Anggraeni, and I. M. Tirta, "Implementasi Random Forest Menggunakan SMOTE untuk Analisis Sentimen Ulasan Aplikasi Sister for Students UNEJ," *J. Nas. Teknol. dan Sist. Inf.*, vol. 9, no. 2, pp. 163–172, 2023, doi: 10.25077/tekno.v9i2.2023.163-172.
- A. P. Natasuwarna, "Analisis Sentimen Keputusan Pemindahan Ibukota Negara Menggunakan Klasifikasi Naive Bayes," *SENSITIF* 2019, pp. 47–53, 2018.
- A. R. Raharja, A. Pramudianto, and Y. Muchsam, "Penerapan Algoritma Decision Tree dalam Klasifikasi Data 'Framingham' Untuk Menunjukkan Risiko Seseorang Terkena Penyakit Jantung dalam 10 Tahun Mendatang".
- A. Setiawan and R. R. Suryono, "Analisis Sentimen Ibu Kota Nusantara menggunakan Algoritma Support Vector Machine dan Naïve Bayes," *Edumatic J. Pendidik. Inform.*, vol. 8, no. 1, pp. 183–192, 2024, doi: 10.29408/edumatic.v8i1.25667.
- B. Santika and A. K. Hidayah, "Implementation of Support Vector Machine for Sentiment Analysis on Bibli App Reviews on Play Store," *J. Vocat. Tek. Elektron. dan Inform.*, 2025.

- D. Aryanti, "Analisis Sentimen Ibukota Negara Baru Menggunakan Metode Naïve Bayes Classifier," *J. Inf. Syst. Res.*, vol. 3, no. 4, pp. 524–531, 2022, doi: 10.47065/josh.v3i4.1944.
- E. Mas'udah, E. D. Wahyuni, and A. A. Arifiyanti, "Analisis Sentimen: Pemindahan Ibu Kota Indonesia Pada Twitter," *J. Inform. dan Sist. Inf.*, vol. 1, no. 2, pp. 397–401, 2020.
- H. Hairani, A. Anggrawan, and D. Priyanto, "Improvement Performance of the Random Forest Method on Unbalanced Diabetes Data Classification Using Smote-Tomek Link," *Int. J. Informatics Vis.*, vol. 7, no. 1, pp. 258–264, 2023, doi: 10.30630/joiv.7.1.1069.
- Hidayah, A. K. (2025). Analisis Sentimen Ulasan Aplikasi Vidio pada Play Store Menggunakan Algoritma Support Vector Machine: Analisis SVM pada ulasan aplikasi vidio. *JUPITER: Jurnal Penelitian Ilmu dan Teknologi Komputer*, 17(1), 147-158.
- M. P. Prapasha and H. Thamrin, "IMPLEMENTASI POS TAGGING DAN ALGORITMA ANTLR PARSER DALAM MEMERIKSA STRUKTUR KALIMAT BAHASA INDONESIA," *Publ. Univ. Muhammadiyah Surakarta*, vol. 16, no. 2, pp. 39–55, 2024.
- N. Istiqamah and M. Rijal, "Klasifikasi Ulasan Konsumen Menggunakan Random Forest dan SMOTE," *J. Syst. Comput. Eng.*, vol. 5, no. 1, pp. 66–77, 2024, doi: 10.61628/jsce.v5i1.1061.
- N. S. Wardani, A. Prahutama, and P. Kartikasari, "Analisis Sentimen Pemindahan Ibu Kota Negara Dengan Klasifikasi Naïve Bayes Untuk Model Bernoulli Dan Multinomial," *J. Gaussian*, vol. 9, no. 3, pp. 237–246, 2020, doi: 10.14710/j.gauss.v9i3.27963.
- P. Arsi, B. A. Kusuma, and A. Nurhakim, "Analisis Sentimen Pindah Ibu Kota Berbasis Naive Bayes Classifier," *J. Inform. Upgris*, vol. 7, no. 1, pp. 1–6, 2021, doi: 10.26877/jiu.v7i1.7636.
- P. Arsi, R. Wahyudi, and R. Waluyo, "Optimasi SVM Berbasis PSO pada Analisis Sentimen Wacana Pindah Ibu Kota Indonesia,"

- J. RESTI, vol. 5, no. 2, pp. 231–237, 2021, doi: 10.29207/resti.v5i2.2698.
- P. Studi, S. Informasi, D. Informatika, F. T. Industri, U. Atma, and J. Yogyakarta, “Analisis Komparatif Algoritme Machine Learning pada Klasifikasi Kualitas Air Layak Minum,” vol. 2, no. 1, pp. 43–55, 2022.
- R. C. Bhagat and S. S. Patil, “Enhanced SMOTE Algorithm for Classification of Imbalanced Big-Data using Random Forest,” IEEE Int. Adv. Comput. Conf., pp. 403–408, 2015.
- S. F. Huwaida, R. Kusumawati, and B. Isnaini, “Analisis Sentimen Komentar YouTube terhadap Pemindahan Ibu Kota Negara Menggunakan Metode Naïve Bayes,” Jambura J. Informatics, vol. 6, no. 1, pp. 26–39, 2024, doi: 10.37905/jji.v6i1.24718.
- S. Sidiq, P. Korespondensi, and N. Shobi Maburur, “Pengembangan Model Prediksi Risiko Diabetes Menggunakan Pendekatan AdaBoost dan Teknik Oversampling SMOTE,” J. Ilm. Inform. Dan Ilmu Komput., vol. 4, pp. 13–23, 2025, [Online]. Available: <https://doi.org/10.58602/jima-ilkom.v4i1.41>
- Syahril Dwi Prasetyo, Shofa Shofiah Hilabi, and Fitri Nurapriani, “Analisis Sentimen Relokasi Ibukota Nusantara Menggunakan Algoritma Naïve Bayes dan KNN,” J. KomtekInfo, vol. 10, pp. 1–7, 2023, doi: 10.35134/komtekinfo.v10i1.330.
- T. Sela and A. Sonita, “Comparison of Naive Bayes and SVM Algorithms in Sentiment Analysis,” *J. Ilm. Tek. Inform. dan Sist. Inf.*, pp. 1–12, 2025.
- U. Findawati and A. M. Rosid, *Buku Ajar Teks Mining*. 2020.



TENTANG PENULIS

Agung Kharisma Hidayah, S.Kom., M.Kom.

Penulis dengan nama panggilan “Agung” ini lahir di Kabupaten Kerinci Propinsi Jambi. Menyelesaikan Sekolah Dasar, Sekolah Menengah Pertama dan Sekolah Menengah Atas di Kabupaten Kerinci tempat kelahirannya. Penulis melanjutkan pendidikan ke Universitas Muhammadiyah Bengkulu Program Studi Teknik Informatika, lulus pada tahun 2015. Pada tahun yang sama Penulis melanjutkan kuliah di Universitas Amikom Yogyakarta, Program Studi S2 Magister Teknik Informatika dan lulus pada tahun 2017. Sejak tahun 2017 hingga sekarang Penulis aktif sebagai dosen di Program studi Teknik Informatika, Universitas Muhammadiyah Bengkulu. Penulis dapat dihubungi melalui Email: kharisma@umb.ac.id

Dandi Sunardi, M.Kom.

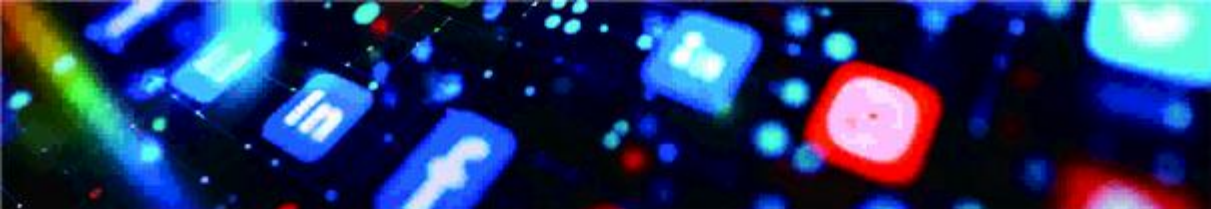
Penulis dengan nama panggilan “Kang Dandi” ini lahir di Kabupaten Kuningan Propinsi Jawa Barat. Menyelesaikan Sekolah Dasar, Sekolah Menengah Pertama di Kabupaten Kuningan tempat kelahiran, Sekolah Menengah Atas di Kabupaten Bengkulu Utara Provinsi Bengkulu. Penulis melanjutkan pendidikan ke Universitas Muhammadiyah Bengkulu, lulus pada tahun 2003. Pada tahun 2015 Penulis melanjutkan kuliah di Universitas Amikom Yogyakarta, Program Studi S2 Magister Teknik Informatika dan lulus pada tahun 2017. Sejak tahun 2017 hingga sekarang Penulis aktif sebagai dosen di Program studi Teknik Informatika, Universitas Muhammadiyah Bengkulu, trainer digital marketing, dan aktif di unit humas, media

dan promosi Universitas Muhammadiyah Bengkulu. Penulis dapat dihubungi melalui Email: dandisunardi@umb.ac.id

Eka Sahputra, S.Kom., M.Kom

Eka Sahputra, lahir di Pasar Baru-Bintuhan pada 20 Mei 1988 dan sekarang berdomisili di Banyumas, Jawa Tengah, menempuh pendidikan dasar hingga menengah di Kabupaten Kaur, Provinsi Bengkulu. Sejak masa sekolah, ia aktif dalam berbagai kegiatan ekstrakurikuler seperti Pencak Silat, Paskibraka, drumband, Pramuka, dan OSIS.

Penulis meraih gelar Sarjana (S-1) dari Program Studi Teknik Informatika, Universitas Muhammadiyah Bengkulu lulus pada tahun 2011, dan melanjutkan pendidikan Pascasarjana (S-2) pada program studi yang sama di Universitas Amikom Yogyakarta hingga lulus tahun 2017. Sejak tahun 2017, penulis mengabdikan sebagai dosen serta aktif dalam kegiatan penelitian pembangunan daerah, pendampingan dan pemberdayaan masyarakat, narasumber Literasi Digital Nasional, dan menjabat sebagai Project Manager dalam proyek Sistem Informasi Manajemen Rumah Sakit (SIMRS) di berbagai daerah di Provinsi Papua dan Nusa Tenggara Barat (NTB) pada tahun 2024–2025. Penulis dapat dihubungi melalui email: eka88edu@gmail.com dan ekasahputra@telkomuniversity.ac.id



ANALISIS SENTIMEN

di Era Digital Konsep, Algoritma,
dan Studi Kasus
Media Sosial

Buku ini hadir sebagai jawaban atas kebutuhan akan pemahaman analisis sentimen di era digital. Dengan judul “Analisis Sentimen di Era Digital: Konsep, Algoritma, dan Studi Kasus Media Sosial”, kami berusaha menggabungkan landasan teoritis dengan pendekatan teknikal yang aplikatif. Pembaca tidak hanya akan diajak memahami konsep dasar seperti sentiment analysis, natural language processing (NLP), dan machine learning, tetapi juga akan diperkenalkan dengan studi-studi kasus nyata yang mencerminkan dinamika sosial di dunia maya.

Kami memilih pemindahan Ibu Kota Negara (IKN) sebagai salah satu studi kasus utama dalam buku ini karena isu tersebut memicu diskusi yang luas dan beragam di masyarakat. Kami juga menyertakan berbagai contoh analisis sentimen lain di platform media sosial dan toko aplikasi, agar pembaca dapat melihat bagaimana teknik yang sama diterapkan pada konteks yang berbeda.

Buku ini ditujukan untuk mahasiswa, dosen, peneliti, praktisi data, dan siapa pun yang ingin memahami atau mengembangkan sistem analisis opini publik secara kuantitatif. Kami menyadari bahwa masih banyak ruang untuk penyempurnaan, baik dari sisi kedalaman bahasan maupun keluasan aplikasinya. Namun, kami berharap buku ini dapat menjadi pijakan awal yang kokoh bagi pembaca untuk mengeksplorasi lebih lanjut potensi besar analisis sentimen di era digital ini.



✉ literasinusantaraofficial@gmail.com
🌐 www.penerbitlitnus.co.id
📖 Literasi Nusantara
📞 literasinusantara_
☎ 085755971589

Pendidikan

+17

