

ANALISIS EMOSI AUDIENS PADA KOMENTAR INSTAGRAM @JTV_REK BERBASIS LATENT SEMANTIC ANALYSIS DENGAN PENDEKATAN ENSEMBLE STACKING UNTUK KLASIFIKASI EMOSI

Jasmine Aulia, Amri Muhaimin, Andri Fauzan Adziima

Sains Data, UPN "VETERAN" JAWA TIMUR

Jalan Rungkut Madya, Gn. Anyar, Kec. Gunung Anyar, Surabaya, Indonesia

22083010074@student.upnjatim.ac.id

ABSTRAK

Perkembangan media sosial telah mengubah pola interaksi masyarakat dalam menyampaikan opini dan ekspresi emosional terhadap berbagai konten digital, khususnya melalui komentar pada platform Instagram. Namun, analisis emosi pada komentar media sosial menghadapi tantangan seperti ketidakseimbangan data, teks yang pendek, dan penggunaan bahasa informal. Penelitian ini bertujuan untuk mengklasifikasikan emosi audiens pada komentar Instagram akun @JTV_rek menggunakan metode Stacking Ensemble dengan dukungan *Latent Semantic Analysis* (LSA) untuk analisis topik. Dataset yang digunakan terdiri dari 3.701 komentar yang dikumpulkan pada periode Januari 2024 hingga Agustus 2025 dan dilabeli secara manual ke dalam tiga kelas emosi, yaitu joy, netral, dan negative. Representasi fitur menggunakan TF-IDF, sedangkan LSA digunakan untuk mengekstraksi pola semantik. Proses klasifikasi dilakukan dengan menggabungkan Random Forest, Support Vector Machine, dan XGBoost dalam skema Stacking Ensemble dengan Logistic Regression sebagai *meta-learner*. Hasil penelitian menunjukkan bahwa model terbaik diperoleh pada skenario Stacking tanpa SMOTE dengan nilai accuracy sebesar 74% dan F1-score sebesar 73%. Hasil ini menunjukkan bahwa pendekatan ensemble mampu meningkatkan performa klasifikasi emosi pada data komentar media sosial.

Kata kunci : Instagram, Klasifikasi Emosi, Latent Semantic Analysis, Stacking Ensemble, TF-IDF

1. PENDAHULUAN

Perkembangan media sosial dalam beberapa tahun terakhir telah mengubah cara masyarakat mengekspresikan opini, pengalaman, dan emosi secara terbuka, khususnya terhadap konten dan industri media [1]. Platform media sosial seperti Instagram dimanfaatkan audiens untuk menyampaikan komentar dan respons publik terhadap konten yang dipublikasikan [2]. Menurut data Statista, jumlah pengguna Instagram di Indonesia mencapai 107,6 juta pada tahun 2025. Temuan ini sejalan dengan laporan DataReportal 2026 yang mencatat sekitar 108 juta audiens iklan Instagram di Indonesia [3]. Interaksi tersebut membentuk ruang diskursif digital yang memungkinkan audiens memberikan evaluasi, kritik, maupun apresiasi secara langsung dan terdokumentasi. Komentar yang ditinggalkan pada unggahan program tidak hanya merepresentasikan bentuk partisipasi digital, tetapi juga mencerminkan respons emosional terhadap isi tayangan [4].

Namun, karakteristik komentar media sosial yang cenderung informal, multitopik, dan mengandung variasi bahasa menimbulkan tantangan dalam proses analisis teks. Permasalahan seperti *noise data*, konteks tersembunyi, dan keterbatasan representasi fitur masih menjadi tantangan utama [5]. Selain itu, data komentar yang bersifat pendek dan tidak terstruktur sering kali menyulitkan model dalam menangkap makna secara kontekstual.

Beberapa penelitian terdahulu menunjukkan bahwa analisis sentimen dan pemodelan topik telah banyak diterapkan untuk memahami opini publik di media sosial. Penelitian oleh [6] menggunakan

pendekatan ensemble membuktikan bahwa penggabungan beberapa algoritma klasifikasi mampu meningkatkan akurasi dibandingkan model tunggal. Lebih lanjut penelitian oleh [7] menganalisis komentar YouTube Mata Najwa menunjukkan bahwa algoritma *Support Vector Machine* (SVM) efektif digunakan dalam memetakan persepsi publik terhadap isu pendidikan tinggi meskipun masih menghadapi tantangan pada teks informal. Sementara itu, pada penelitian [8] menunjukkan bahwa metode *Latent Semantic Analysis* (LSA) efektif dalam merepresentasikan struktur semantik laten dan melakukan pemodelan topik pada 720 abstrak skripsi dengan menghasilkan 37 topik serta tingkat akurasi rata-rata sebesar 79% presisi hingga 80%, dan recall mencapai 90%, sehingga membuktikan kemampuan LSA dalam mengekstraksi pola makna teks secara sistematis dan akurat.

Meskipun demikian, sebagian besar penelitian masih berfokus pada analisis sentimen secara umum dan belum banyak yang mengintegrasikan pendekatan pemodelan topik dengan metode ensemble secara bersamaan untuk meningkatkan akurasi dan stabilitas model. Selain itu, penerapan metode tersebut pada konteks komentar media sosial lokal, khususnya pada industri penyiaran, masih terbatas. Maka dari itu, penelitian ini mengusulkan pendekatan *Stacking Ensemble* dengan dukungan *Latent Semantic Analysis* (LSA) untuk mengklasifikasikan emosi audiens pada komentar Instagram @JTV_REK. Pendekatan ini bertujuan untuk meningkatkan performa klasifikasi emosi serta mengidentifikasi pola emosi yang muncul dalam interaksi audiens, sehingga dapat memberikan

kontribusi dalam pengembangan metode analisis emosi berbasis teks serta pemanfaatannya dalam evaluasi konten media digital.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen dan Emosi Teks

Analisis sentimen merupakan proses identifikasi dan klasifikasi opini dalam teks ke dalam kategori tertentu, seperti positif, negatif, dan netral. Dalam perkembangannya, analisis tidak hanya terbatas pada polaritas sentimen, tetapi juga pada identifikasi emosi spesifik seperti senang (*joy*), marah (*anger*), sedih (*sadness*), dan takut (*fear*) [9]. Analisis emosi memberikan informasi yang lebih granular dibandingkan analisis sentimen karena mampu menangkap kondisi afektif yang lebih kompleks dalam teks. Pada media sosial, analisis emosi menjadi lebih menantang karena karakteristik teks yang informal, penggunaan slang, emoji, serta campuran bahasa. Oleh karena itu, pendekatan berbasis machine learning dan representasi semantik diperlukan untuk mengklasifikasikan emosi secara sistematis.

2.2. Natural Language Processing (NLP)

Natural Language Processing (NLP) merupakan cabang kecerdasan buatan yang berfokus pada pemrosesan dan analisis bahasa alami oleh komputer [10]. Dalam klasifikasi teks, NLP berperan dalam mengubah teks tidak terstruktur menjadi representasi numerik yang dapat diproses oleh algoritma machine learning.

2.3. Latent Semantic Analysis (LSA)

Latent Semantic Analysis atau LSA merupakan metode statistik yang digunakan untuk mengekstraksi dan memrepresentasikan makna kata berdasarkan konteks kemunculkannya dalam suatu korpus teks [11]. LSA memodelkan hubungan antar kata dalam bentuk matriks semantik yang kemudian diproses menggunakan *Singular Value Decomposition* (SVD) untuk mengidentifikasi dimensi laten paling signifikan serta melakukan reduksi dimensi. Metode ini bersifat statistik murni dan tidak bergantung pada struktur linguistik, jaringan semantik, kamus, maupun basis pengetahuan, melainkan hanya memanfaatkan teks mentah yang diuraikan menjadi unit kata dan dianalisis pada tingkat kalimat atau paragraf.

2.4. Ensemble Learning

Ensemble learning merupakan pendekatan machine learning yang menggabungkan beberapa model prediktif untuk menghasilkan kinerja yang lebih baik dibandingkan model Tunggal [12]. Metode ini bekerja dengan memanfaatkan keberagaman model sehingga kesalahan dari satu model dapat dikompensasi oleh model lainnya. Pendekatan ensemble meliputi teknik seperti bagging, boosting, voting dan stacking untuk meningkatkan performa model [13].

2.5. Algoritma Base Learner

Base learner merupakan model dasar yang digunakan dalam kerangka ensemble learning untuk menghasilkan prediksi awal sebelum dikombinasikan. Keberagaman karakteristik untuk menghasilkan prediksi awal sebelum dikombinasikan. Keberagaman karakteristik antar base learner menjadi faktor penting dalam meningkatkan kinerja ensemble, karena perbedaan pola kesalahan antar model dapat saling melengkapi dan memperbaiki kemampuan generalisasi sistem.

2.6. Random Forest

Random Forest merupakan metode klasifikasi berbasis ensemble yang terdiri dari sejumlah pohon keputusan (*decision tree*) yang dibangun menggunakan pemilihan fitur dan sampel data secara acak untuk meningkatkan akurasi serta mengurangi *overfitting*. Setiap pohon menghasilkan prediksi yang kemudian digabungkan untuk menentukan hasil akhir klasifikasi [14]. Prediksi klasifikasi dirumuskan sebagai:

$$\hat{y} = \text{mode}\{h_b(x)\}_{b=1}^B \quad (1)$$

Berdasarkan persamaan (1), $h_b(x)$ adalah prediksi dari pohon ke- b , dan B adalah jumlah pohon.

2.7. Support Vector Machine (SVM)

Support Vector Machine atau SVM merupakan algoritma pembelajaran bertujuan menemukan *hyperplane* optimal yang mampu memisahkan data antar kelas dengan margin maksimum. SVM bekerja dengan memilih beberapa titik data yang disebut *support vectors* sebagai penentu batas keputusan. Untuk menangani data nonlinier, SVM dapat menggunakan fungsi kernel guna memetakan data ke ruang berdimensi lebih tinggi. Dalam penelitian ini, SVM digunakan sebagai salah satu *base learner* pada metode stacking untuk menghasilkan prediksi awal yang kemudian dikombinasikan dengan model lainnya. Fungsi keputusan SVM dirumuskan sebagai berikut:

$$f(x) = w^T x + b \quad (2)$$

Berdasarkan persamaan (2), w adalah vektor bobot dan b adalah bias.

2.8. Extreme Gradient Boosting (XGBOOST)

Extreme Gradient Boosting (XGBoost) merupakan algoritma berbasis *boosting* yang membangun model secara bertahap dengan mengoptimalkan fungsi kerugian menggunakan pendekatan gradien. Pada setiap iterasi, model baru dibentuk untuk memperbaiki kesalahan prediksi model sebelumnya. XGBoost juga menerapkan teknik regularisasi untuk mengontrol kompleksitas model sehingga mampu mengurangi risiko *overfitting* dan meningkatkan kemampuan generalisasi. Model prediksi XGBoost dinyatakan dalam bentuk aditif sebagai berikut:

$$\hat{y}_i = \sum_{k=1}^K f_k(x_i) \quad (3)$$

Dalam persamaan (3), f_k adalah fungsi pohon keputusan ke- k , dan K adalah jumlah pohon yang dibangun secara bertahap.

2.9. Penelitian Terdahulu

Penelitian mengenai analisis sentimen pada media sosial telah banyak dilakukan dengan berbagai pendekatan klasifikasi. Studi tentang komentar Instagram terkait polemik desain jersey tim nasional Indonesia menunjukkan bahwa penggunaan model ensemble berbasis *VotingClassifier* yang menggabungkan Random Forest, Naive Bayes, dan Support Vector Machine mampu meningkatkan akurasi hingga 77,4% lebih tinggi dibandingkan model tunggal [6]. Temuan ini menegaskan bahwa pendekatan ensemble efektif dalam meningkatkan stabilitas dan performa klasifikasi pada data media sosial yang tidak terstruktur.

Penelitian sebelumnya yang menganalisis komentar YouTube Mata Najwa menggunakan algoritma *Support Vector Machine* (SVM) menunjukkan bahwa SVM mampu memetakan persepsi publik terhadap isu pendidikan tinggi, meskipun masih menghadapi tantangan dalam menangani teks informal berbahasa Indonesia [7]. Temuan tersebut menegaskan bahwa SVM tetap relevan dalam klasifikasi teks media sosial, khususnya pada data berdimensi tinggi. Sementara itu, penelitian berbasis *Latent Semantic Analysis* (LSA) pada pemodelan topik abstrak skripsi menunjukkan bahwa metode ini mampu mengekstraksi struktur semantik laten secara sistematis dengan tingkat akurasi rata-rata sebesar 79%, presisi hingga 80%, dan recall mencapai 90% [8], sehingga membuktikan efektivitas LSA dalam merepresentasikan makna tersembunyi dalam teks dan meningkatkan kualitas fitur klasifikasi. Meskipun demikian, integrasi representasi semantik berbasis LSA dengan pendekatan *stacking ensemble* dalam konteks analisis emosi audiens program televisi lokal pada komentar Instagram masih relatif terbatas. Oleh karena itu, penelitian ini mengusulkan penggabungan LSA dengan skema *multitopic stacking ensemble* untuk meningkatkan stabilitas dan akurasi klasifikasi emosi pada data komentar Instagram @JTV_REK.

3. METODE PENELITIAN



Gambar 1. Alur penelitian

3.1. Data Collection

Data yang digunakan dalam penelitian ini diperoleh melalui proses *scraping* komentar pada akun Instagram @JTV_rek untuk periode Januari 2024

hingga Agustus 2025. Proses pengambilan data dilakukan menggunakan teknik *scraping* dengan bantuan perangkat lunak atau library pemrograman yang mendukung ekstraksi data dari media sosial. Data yang dikumpulkan berupa teks komentar pengguna yang kemudian disimpan dalam format terstruktur untuk keperluan analisis lebih lanjut. Pemilihan akun dan periode waktu tersebut didasarkan pada relevansi topik serta ketersediaan data yang representatif terhadap tujuan penelitian.

3.2. Data Preprocessing

Tahap data preprocessing dilakukan untuk membersihkan dan menyiapkan data komentar sebelum digunakan dalam proses analisis dan pemodelan. Teks yang diperoleh cenderung mengandung noise seperti emoji, singkatan, serta karakter tidak baku. Oleh karena itu, dilakukan beberapa tahapan sebagai berikut:

- Data Cleaning**
Menghapus karakter khusus, tanda baca, angka, URL, emoji, serta komentar duplikat agar data menjadi lebih terstruktur.
- Case Folding**
Mengubah seluruh teks menjadi huruf kecil untuk menjaga konsistensi format.
- Normalisasi**
Mengubah kata tidak baku atau slang menjadi bentuk baku agar diproses oleh model.
- Stopword Removal**
Menghapus kata-kata umum yang tidak memiliki makna signifikan dalam analisis, seperti kata penghubung dan kata umum lainnya.
- Emoticon Handling**
Mempertahankan atau mengolah emoticon sebagai bagian dari ekspresi emosi dalam teks.

3.3. Feature Extraction

Setelah data melalui tahap preprocessing, Langkah selanjutnya adalah mengubah teks menjadi representasi numerik menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*). Metode ini menghitung bobot setiap kata berdasarkan frekuensi kemunculannya dalam suatu dokumen serta tingkat kepentingannya dalam keseluruhan korpus. Hasil dari proses ini berupa matriks fitur yang digunakan sebagai input dalam proses pelatihan model machine learning.

3.4. Modelling

Tahap modelling dilakukan menggunakan pendekatan *stacking ensemble*. Proses ini terdiri dari dua tahap utama, yaitu pelatihan *base learner* dan pembentukan model *stacking* menggunakan *meta-learner*.

3.5. Base Learners

Tiga algoritma digunakan sebagai *base learners*, yaitu Random Forest, Support Vector Machine (SVM), dan Extreme Gradient Boosting (XGBoost).

Masing-masing model dilatih menggunakan data latih untuk menghasilkan prediksi awal terhadap data.

3.6. Stacking Ensemble

Metode stacking dilakukan dengan menggabungkan hasil prediksi dari ketiga *base learners* sebagai fitur baru. Output prediksi tersebut kemudian digunakan sebagai input bagi *meta-learner* untuk menghasilkan prediksi akhir. Secara matematis, proses stacking dapat dinyatakan sebagai:

$$z_i = h_1(x_i), h_2(x_i), h_3(x_i) \quad (4)$$

$$\hat{y}_i = g(z_i) \quad (5)$$

Berdasarkan persamaan (4) dan (5), $h_1(x_i), h_2(x_i), h_3(x_i)$ adalah *base learners* (RF, SVM, dan XGBoost), g adalah *meta-learner*, dan $\hat{y}_i$ adalah prediksi akhir.

3.7. Insight dan Evaluasi

Tahap akhir penelitian adalah evaluasi dan interpretasi hasil model klasifikasi berbasis *stacking ensemble*. Evaluasi dilakukan menggunakan *confusion matrix*, yang terdiri dari empat komponen utama, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Berdasarkan *confusion matrix* tersebut, dihitung beberapa metrik evaluasi sebagai berikut:

a. Accuracy

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (6)$$

Accuracy menunjukkan proporsi keseluruhan prediksi yang benar dibandingkan total data yang diuji. Nilai accuracy yang tinggi mengindikasikan bahwa model memiliki performa klasifikasi yang baik secara umum.

b. Precision

$$Precision = \frac{TP}{TP+FP} \quad (7)$$

Precision mengukur tingkat ketepatan model dalam memprediksi kelas positif. Nilai precision yang tinggi menunjukkan bahwa sebagian besar prediksi positif yang dihasilkan model benar dan minim kesalahan tipe false positive.

c. Recall

$$Recall = \frac{TP}{TP+FN} \quad (8)$$

Recall mengukur kemampuan model dalam mendeteksi seluruh data yang benar-benar termasuk dalam kelas positif. Nilai recall yang tinggi menunjukkan bahwa model mampu menangkap sebagian besar data positif dan meminimalkan kesalahan tipe false negative.

d. F1-Score

$$F1 = 2 \times \frac{Precision \times Recall}{Precision+Recall} \quad (9)$$

F1-score merupakan rata-rata harmonis antara *precision* dan *recall*. Metrik ini digunakan untuk memberikan gambaran keseimbangan antara ketepatan dan kelengkapan prediksi model, terutama ketika distribusi kelas tidak seimbang.

4. HASIL DAN PEMBAHASAN

4.1. Data Collecting

Data penelitian dikumpulkan dari platform Instagram pada akun resmi @jtv_rek dalam periode Januari 2024 hingga Agustus 2025. Proses pengambilan data dilakukan dengan mengekstraksi komentar publik pada unggahan menggunakan ekstensi peramban *IG Comments Export Tool* dengan memasukkan URL setiap postingan ke dalam sistem dan menjalankan fitur *Start Export*. Hanya postingan yang memiliki komentar, sehingga dataset yang diperoleh benar-benar merepresentasikan interaksi audiens. Hasil ekstraksi kemudian disimpan dalam format CSV untuk selanjutnya melalui tahap pembersihan data dan pemrosesan lebih lanjut. Data yang digunakan berupa komentar publik tanpa memuat informasi pribadi pengguna secara spesifik, sehingga tetap berada dalam batas penggunaan data terbuka untuk kepentingan penelitian akademik.

Tabel 1. Data Collecting

text
Seger e min cocok wes udan udan ngene
@stasiunandangdutyv Tiket nonton film Ambyar Mak Byare kekno aku wae min, emakku pengen ndelok pilm e Happy Asmara-Gilga iki ket iko,
Seru bgt💖
Love you JTV, terimakasih telah menerima kami💖

Berdasarkan Tabel 1, terlihat bahwa data komentar yang digunakan dalam penelitian memiliki karakteristik bahasa yang bersifat informal, singkat, dan beragam, termasuk penggunaan bahasa daerah serta ekspresi emosional seperti pujian dan apresiasi. Komentar yang ditampilkan menunjukkan adanya variasi bentuk penyampaian opini, mulai dari ungkapan sederhana hingga kalimat yang lebih panjang dan kontekstual. Hal ini mencerminkan kondisi nyata interaksi pengguna di media sosial, yang cenderung tidak terstruktur dan mengandung ambiguitas, sehingga menjadi tantangan tersendiri dalam proses analisis dan klasifikasi emosi.

4.2. Data Preprocessing

Tahap preprocessing dilakukan untuk membersihkan dan menyiapkan data komentar sebelum proses ekstraksi fitur. Karakteristik komentar media sosial yang tidak terstruktur dan mengandung bahasa tidak baku, dilakukan pembersihan awal dengan menghapus karakter khusus, tanda baca, serta emoji. Selanjutnya dilakukan normalisasi untuk mengubah kata tidak baku dan singkatan menjadi bentuk yang lebih formal.

Tabel 2. Proses Data Preprocessing

Tahap	Teks
Teks asli	Wuiiiih segerrrr rek 😊
Casefold	wuiiiih segerrrr rek 😊
Cleaning	wuiiiih segerrrr rek suka
Normalisasi	wuih segar rek suka

Tahap	Teks
Tokenizing	[wuih, segar, rek, suka]
Stopword Removal	[wuih, segar, rek, suka]
Final_text	wuih segar rek suka

Berdasarkan Tabel 2, proses preprocessing dilakukan secara bertahap untuk mengubah teks mentah menjadi bentuk yang lebih terstruktur dan siap diproses oleh model. Tahap casefold bertujuan menyeragamkan huruf, sedangkan emoticon pada teks dikonversi menjadi kata yang merepresentasikan emosi. Selanjutnya, normalisasi memperbaiki kata tidak baku menjadi bentuk standar, diikuti dengan tokenizing untuk memecah teks menjadi unit kata. Proses stopwords removal dilakukan untuk menghilangkan kata yang kurang bermakna, sehingga diperoleh representasi akhir teks yang lebih bersih dan informatif, yaitu “wuih segar rek suka”.

4.3. Feature Extraction

Setelah melalui tahap *preprocessing*, data teks kemudian diubah menjadi representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Representasi TF-IDF yang dihasilkan kemudian digunakan sebagai input dalam proses reduksi dimensi menggunakan metode *Latent Semantic Analysis* (LSA) yang bertujuan untuk mengekstraksi pola semantik laten serta mengelompokkan kata-kata ke dalam beberapa topik berdasarkan kemiripan konteks kemunculannya. Namun, dalam eksperimen ini, TF-IDF menunjukkan performa yang lebih optimal dibandingkan representasi berbasis LSA dalam proses klasifikasi.

Tabel 3. Evaluasi Jumlah Topik LSA

k Topik	Explained variance
1	0.0220
2	0.0391
3	0.0513
4	0.0627
5	0.0710
6	0.0780
7	0.0830
8	0.0870
9	0.0900
10	0.0920
11	0.0940
12	0.0960
13	0.0980
14	0.0990
15	0.1000
16	0.1440
17	0.1486
18	0.1530
19	0.1576
20	0.1618

Tabel 3 menampilkan jumlah topik pada metode LSA diuji dengan variasi nilai k dari 1 hingga 20. Evaluasi dilakukan menggunakan nilai explained variance untuk melihat seberapa besar informasi yang dapat dipertahankan setelah reduksi dimensi. Berdasarkan hasil pengujian, nilai explained variance cenderung meningkat seiring dengan bertambahnya jumlah topik. Berdasarkan hasil evaluasi, dipilih nilai k terbaik yang digunakan pada tahap selanjutnya. Pemilihan nilai k ini mempertimbangkan keseimbangan antara kompleksitas model dan kemampuan representasi fitur.

Tabel 4. Hasil Topik LSA

Topik	Kata Dominan	Bobot
Topik 1	tertawa, tidak, di, suka, ini, yang, cinta	0.9740, 0.0703, 0.0672, 0.0630, 0.0447, 0.0436, 0.0399
Topik 2	suka, cinta, keren, bu, risma, jtv, kasih	0.8871, 0.3319, 0.1123, 0.0985, 0.0780, 0.0740, 0.0478
Topik 3	cinta, bu, risma, yang, tidak, ya, 01	0.6740, 0.3488, 0.3269, 0.0901, 0.0886, 0.0729, 0.0723
Topik 4	bu, risma, luluk, tidak, 01, jatim, menyala	0.5130, 0.4479, 0.1145, 0.1085, 0.1022, 0.0949, 0.0851
Topik 5	risma, bu, cinta, tertawa, suka, hans, ini	0.3054, 0.2893, 0.1557, 0.1257, 0.0988, 0.0212, 0.0165

Berdasarkan Tabel 4, hasil pemodelan LSA menunjukkan bahwa komentar audiens didominasi oleh kata-kata yang mencerminkan ekspresi positif, seperti “suka”, “cinta”, “keren”, dan “tertawa”. Selain itu, muncul pula kata yang berkaitan dengan tokoh atau konteks tertentu, seperti “Risma” dan “Luluk”, yang menunjukkan adanya pembahasan mengenai figur publik dalam komentar. Secara umum, topik-topik yang terbentuk mengindikasikan bahwa audiens memberikan respon yang cenderung positif serta interaktif terhadap konten yang dipublikasikan. Beberapa kata umum seperti “di”, “yang”, dan “tidak” masih muncul karena keterbatasan proses preprocessing, namun tidak mempengaruhi secara signifikan dalam interpretasi topik. Representasi topik yang dihasilkan oleh LSA selanjutnya digunakan sebagai fitur dalam proses klasifikasi menggunakan model machine learning.

4.4. Emotion Labelling

Proses pelabelan emosi dilakukan secara manual dengan mengkategorikan setiap komentar ke dalam salah satu kelas emosi yang telah ditentukan. Penentuan label emosi dilakukan dengan menginterpretasikan makna dan konteks dari setiap komentar yang ditulis oleh pengguna. Setiap komentar kemudian diklasifikasikan ke dalam lima kategori emosi, yaitu *joy*, *anger*, *sadness*, *fear*, dan *netral*.

Tabel 5. Labeling Emosi

Emotion_label	Count
Netral	1103
Joy	1686
Anger	252
Sadness	608
Fear	52

Berdasarkan Tabel 5, distribusi label emosi menunjukkan ketidakseimbangan yang cukup signifikan, dimana kelas *joy* mendominasi dengan jumlah 1686 data, diikuti oleh *netral* sebanyak 1103 data. Sementara itu, kelas *sadness* dan *anger* memiliki jumlah yang lebih rendah, dan kelas *fear* merupakan yang paling sedikit dengan hanya 52 data.

Ketimpangan ini menunjukkan bahwa data cenderung didominasi oleh emosi positif dan netral, sehingga berpotensi memengaruhi kinerja model dalam mengenali emosi minoritas secara optimal.

Tabel 6. *Class distribution*

Emotion_label	Count
Netral	1103
Joy	1686
Negative	912

Untuk mengatasi ketidakseimbangan distribusi kelas, dilakukan proses konsolidasi yang ditunjukkan pada Tabel 6. Kelas *joy* tetap memiliki jumlah tertinggi sebanyak 1633 data, diikuti oleh *netral* sebanyak 1090 data, dan *negative* sebanyak 908 data yang merupakan gabungan dari *anger*, *sadness*, dan *fear*. Pengelompokan ini berhasil mengurangi tingkat ketimpangan antar kelas, sehingga diharapkan dapat membantu model dalam mempelajari pola emosi secara lebih optimal dan meningkatkan kinerja klasifikasi secara keseluruhan.

4.5. Modelling

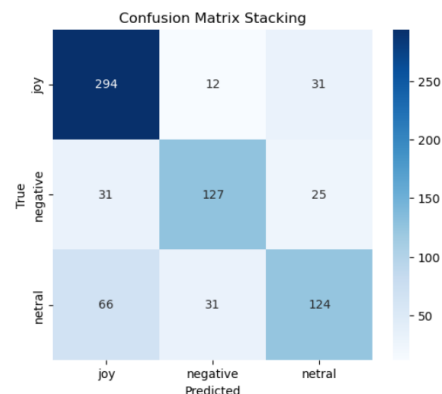
Pada tahap ini, dilakukan proses klasifikasi emosi menggunakan beberapa algoritma machine learning, yaitu Random Forest, Support Vector Machine (SVM), XGBoost, serta metode Stacking Ensemble sebagai model gabungan. Sebelum dilakukan pemodelan, data yang telah melalui proses *preprocessing* dan ekstraksi fitur menggunakan TF-IDF serta reduksi dimensi menggunakan LSA dibagi menjadi data latih dan data uji dengan perbandingan 80:20. Pembagian data dilakukan menggunakan teknik *stratified sampling* untuk menjaga proporsi distribusi kelas. Evaluasi model dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score* untuk menilai performa model pada setiap kelas emosi secara seimbang. Berdasarkan hasil eksperimen, diperoleh performa model terbaik seperti yang ditunjukkan pada Tabel 7.

Tabel 7. Hasil Performa Model Terbaik

Model	Skenario	Acc	Precision	Recall	F1-score
Stacking	Tanpa SMOTE	0.7395	0.7358	0.7395	0.7340
SVM	Tanpa SMOTE	0.7368	0.7335	0.7368	0.7297
Random Forest	Dengan SMOTE	0.7179	0.7151	0.7179	0.7130
XGBoost	Tanpa SMOTE	0.7017	0.7013	0.7017	0.6923

Berdasarkan Tabel 7, model Stacking Ensemble tanpa SMOTE menunjukkan performa terbaik dibandingkan model lainnya dengan nilai *accuracy* sebesar 0.7395 dan *F1-score* sebesar 0.7340. Hal ini menunjukkan bahwa kombinasi beberapa model dasar mampu meningkatkan kemampuan klasifikasi dibandingkan model tunggal. Model SVM juga menunjukkan performa yang kompetitif dengan nilai akurasi yang mendekati stacking, sedangkan Random

Forest dan XGBoost memiliki performa yang relatif lebih rendah.

Gambar 2. *Confusion Matrix Stacking Ensemble*

Berdasarkan Gambar 2, model *Stacking Ensemble* menunjukkan performa yang baik dalam mengklasifikasikan emosi, yang terlihat dari dominasi nilai pada diagonal confusion matrix. Kelas *joy* memiliki jumlah prediksi benar tertinggi, sementara kelas *negative* dan *netral* juga menunjukkan performa yang cukup stabil. Namun, masih terdapat kesalahan klasifikasi, terutama antara kelas *netral* dan *joy*, yang mengindikasikan adanya kemiripan karakteristik pada kedua kelas tersebut sehingga lebih sulit dibedakan oleh model.

Tabel 8. Hasil Evaluasi Klasifikasi Model

Emosi	Precision	Recall	F1-score	Support
Joy	0.75	0.87	0.81	337
Negative	0.75	0.69	0.72	183
Netral	0.69	0.56	0.62	221
Accuracy			0.74	741

Berdasarkan Tabel 8, kelas *joy* memiliki performa terbaik dengan nilai *recall* tertinggi sebesar 0.87 dan *F1-score* sebesar 0.81, yang menunjukkan bahwa model mampu mengenali emosi positif dengan sangat baik. Kelas *negative* menunjukkan performa yang cukup stabil, meskipun nilai *recall* lebih rendah dibandingkan kelas *joy*. Sementara itu, kelas *netral* memiliki performa paling rendah, khususnya pada nilai *recall* sebesar 0.56, yang menunjukkan bahwa model masih mengalami kesulitan dalam mengidentifikasi emosi netral secara akurat. Hal ini mengindikasikan bahwa emosi netral cenderung lebih ambigu dibandingkan emosi lainnya.

4.6. Evaluasi dan Pembahasan

Berdasarkan hasil pengujian, model *Stacking Ensemble* tanpa SMOTE menunjukkan performa terbaik dengan nilai *accuracy* sebesar 0.7395 dan *F1-score* sebesar 0.7340, yang mengindikasikan bahwa pendekatan ensemble mampu meningkatkan kemampuan klasifikasi dibandingkan model tunggal. Model memiliki performa terbaik pada kelas *joy*, sedangkan kelas *netral* masih sulit diklasifikasikan secara akurat karena sifatnya yang lebih ambigu, yang

juga terlihat dari kesalahan klasifikasi pada confusion matrix terutama antara kelas *netral* dan *joy*. Selain itu, penggunaan TF-IDF terbukti efektif dalam merepresentasikan data teks, sementara penerapan LSA tidak selalu meningkatkan performa karena adanya kemungkinan kehilangan informasi akibat reduksi dimensi. Secara keseluruhan, model cenderung lebih baik dalam mengenali emosi yang bersifat eksplisit dibandingkan emosi yang ambigu.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil menerapkan pendekatan klasifikasi emosi pada komentar Instagram *@jtv_rek* menggunakan metode *Stacking Ensemble* dengan representasi fitur berbasis TF-IDF serta eksplorasi *Latent Semantic Analysis* (LSA) untuk analisis topik. Proses penelitian meliputi tahap *preprocessing* teks, pengelompokan emosi menjadi tiga kelas yaitu *joy*, *negative*, dan *netral*, ekstraksi fitur menggunakan TF-IDF, serta klasifikasi menggunakan metode *Stacking Ensemble* yang menggabungkan Random Forest, Support Vector Machine, dan XGBoost dengan Logistic Regression sebagai *meta-learner*. Berdasarkan hasil evaluasi, model terbaik diperoleh pada skenario Stacking tanpa SMOTE dengan nilai *accuracy* sebesar 0.7395 dan *F1-score* sebesar 0.7340, yang menunjukkan bahwa pendekatan ensemble mampu meningkatkan performa klasifikasi dibandingkan model tunggal. Model memiliki performa terbaik dalam mengenali emosi *joy*, sedangkan kelas *netral* masih menjadi tantangan karena sifatnya yang lebih ambigu dan sering mengalami kesalahan klasifikasi. Selain itu, penggunaan TF-IDF terbukti lebih efektif dibandingkan LSA dalam konteks klasifikasi, meskipun LSA tetap bermanfaat dalam mengekstraksi pola topik dari data komentar. Untuk penelitian selanjutnya, disarankan untuk meningkatkan kualitas dan keseimbangan data, khususnya pada kelas yang sulit diklasifikasikan seperti *netral*, serta mengeksplorasi metode representasi teks berbasis *deep learning* seperti BERT atau IndoBERT untuk memperoleh performa klasifikasi yang lebih optimal.

DAFTAR PUSTAKA

- [1] J. Kaur, A. Sethi, and M. Mittal, "Sentiment Analysis for Predictive Insights in the Media & Entertainment Industry Using Big Data," *South Eastern European Journal of Public Health (SEEJPH)*, vol. XXVI, no. S2, 2025..
- [2] A. Prihanum and D. Fadillah, "Analisis Sentimen Tanggapan Publik di Media Sosial Instagram Terhadap Program CSR 'ESG Existence (EXIST)' PT Telkom," *Jurnal Audiens*, vol. 5, no. 4, 2024.
- [3] Iniloh.com, "Gila! Indonesia Kunci Peringkat 4 Pengguna Instagram Terbanyak Dunia, Terbesar di Asia." [Online]. Available: [https://www.inilah.com/gila-indonesia-kunci-](https://www.inilah.com/gila-indonesia-kunci-peringkat-4-pengguna-instagram-terbanyak-dunia-terbesar-di-asia)
- peringkat-4-pengguna-instagram-terbanyak-dunia-terbesar-di-asia
- [4] O. K. D. Arwansyah, P. Suroso, and H. Sigalingging, "Politik dalam Bingkai Media : Studi Mediatization of Politics terhadap Konten Sosial Media Dedi Mulyadi Politics in the Media Frame : A Study of Mediatization of Politics on Dedi Mulyadi 's Social Media Content," *Journal of Education, Humaniora and Social Sciences (JEHSS)*, vol. 8, no. 1, pp. 190–197, 2025.
- [5] G. D. Devi and S. Kamalakkannan, "Literature Review on Sentiment Analysis in Social Media : Open Challenges toward Applications," *International Journal of Advanced Science and Technology*, vol. 29, no. 7, pp. 1462–1471, 2020.
- [6] I. Pratama, "Analisis Sentimen Komentar Instagram Terhadap Polemik Desain Jersey TIMNAS Indonesia Dengan Metode Ensemble," *Jurnal Multidisiplin Inovatif*, vol. 9, no. 7, pp. 479–488, 2025.
- [7] M. H. Maulana, R. A. Sanchia, and D. A. Prasetya, "Persepsi Publik terhadap Pendidikan Tinggi dalam Komentar YouTube Mata Najwa : Analisis Sentimen dengan Algoritma Support Vector Machine," *Seminar Nasional AMIKOM Surakarta (SEMNAS)*, November, pp. 197–207, 2025.
- [8] R. Hakim, "Topic Modelling Dokumen Skripsi Menggunakan Metode *Latent Semantic Analysis*," skripsi, Program Studi Sistem Informasi, UIN Sunan Ampel Surabaya, 2020.
- [9] E. A. Pusvita, D. Stefanus, P. Ranggup, and U. Arfan, "Machine Learning-Based Sentiment Analysis of Public Opinion on Palm Oil Plantation Expansion in Papua," *Institut Riset dan Publikasi Indonesia (IRPI)*, 2025.
- [10] IBM, "Natural Language Processing." [Online]. Available: <https://www.ibm.com/id-id/think/topics/natural-language-processing>
- [11] A. Kemal Al Izzi and R. A. Pratama, "Komparasi Algoritma Topic Modelling LDA VS LSA Pada Berita Detikcom," *FORMAT: Jurnal Ilmiah Teknik Informatika*, vol. 13, no. 1, pp. 44–54, 2024.
- [12] S. Damayanti, "Implementasi Algoritma Ensemble Learning Untuk Peningkatan Akurasi Prediksi," *Jurnal Dunia Data*, vol. 1, no. 5, pp. 1–15, 2024.
- [13] L. Berriche and A. Loulizi, "Prediction of Failures in the Project Management Knowledge Areas Using Optimized Ensemble Models in Software Companies," *Discover Applied Sciences*, vol. 7, no. 7, 2025, Art. no. 718..
- [14] R. R. Fitri and W. Apriandari, "Penggunaan Random Forest dalam Sistem Klasifikasi Kecemasan pada Generasi Z," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 3, pp. 561–571, 2025.