

SENTIMEN PUBLIK TERHADAP KONTROVERSI BYON COMBAT SHOWBIZ BERBASIS SUPPORT VECTOR MACHINE DAN NAIVE BAYES

Moch. Akbar Ramdani¹⁾, Mohammad Nur Fitriyadi²⁾, Mamok Andri Senubekti³⁾

1. Informatika, Universitas Teknologi Digital
email: moch20122010@digitechuniversity.ac.id
2. Informatika, Universitas Teknologi Digital
email: mohammadnur@digitechuniversity.ac.id
3. Manajemen Informatika, Universitas Teknologi Digital
email: mamokandri@digitechuniversity.ac.id

Abstract

The controversial match result of the sportainment event Byon Combat Showbiz Vol. 6 sparked massive debates on social media. The high volume of comments filled with slang and combat sports jargon makes public opinion mapping prone to subjective bias and Out-Of-Vocabulary (OOV) problems. This study aims to objectively classify netizens' sentiment polarity and compare the performance of Naïve Bayes and Support Vector Machine (SVM) algorithms. The extracted dataset of 1,081 documents was processed using Custom Dictionary-based Normalization to reduce linguistic noise. Features were extracted using TF-IDF weighting, while class imbalance was handled using the Synthetic Minority Over-sampling Technique (SMOTE) strictly on the training data to prevent data leakage. The results showed that public opinion was dominated by negative sentiments at 64.5%, rooted in criticism of the referee's technical regulations. Model evaluation proved that linear kernel-based SVM had the most optimal performance with an accuracy rate of 59.45%, outperforming Naïve Bayes which reached 58.99%. The use of a custom dictionary and strict test data separation in SVM proved effective in precisely mapping public sentiment on a complex syntactic dataset.

Kata Kunci : Byon Combat, Custom Dictionary, Naïve Bayes, Sentiment Analysis, Support Vector Machine

A. PENDAHULUAN

Dinamika media sosial saat ini memegang peranan esensial dalam membentuk opini publik [1]. Selain Instagram, eskalasi kontroversi hasil pertandingan ini juga teramplifikasi secara masif melalui platform TikTok dan YouTube. Karakteristik TikTok yang berbasis video pendek mempercepat

penyebaran cuplikan laga, sementara YouTube menyediakan ruang bagi analisis video mendalam, di mana keduanya memicu akumulasi komentar spontan lintas negara yang memperluas spektrum polaritas opini publik. Momentum ini dimanfaatkan secara strategis oleh industri olahraga kombat melalui konsep *sportainment*, yaitu

penggabungan kompetisi adu fisik dengan elemen hiburan teatrikal [2].

Implementasi konsep tersebut terlihat jelas pada perhelatan Byon Combat Showbiz volume enam. Keputusan medis yang menghentikan partai utama antara Ronal Siahaan dan Putra Abdullah memicu kontroversi berskala besar di ekosistem digital [3]. Penghentian sepihak yang menggugurkan potensi ronde tambahan (*judgement round*) tersebut dinilai mencederai ekspektasi audiens, sehingga memicu lonjakan komentar warganet. Volume data tekstual yang masif membuat pemetaan opini secara manual menjadi sangat tidak efisien dan rentan terhadap bias subjektif. Pendekatan klasifikasi dalam arsitektur komputasi secara konsisten diaplikasikan pada berbagai domain komparasi data [4], termasuk dalam pemetaan opini tekstual dari arsip data yang masif. Oleh karena itu, pendekatan komputasional melalui *text mining* dan analisis sentimen diimplementasikan guna mengekstraksi opini acak tersebut menjadi data kuantitatif yang objektif [5].

Penelitian ini dikembangkan berdasarkan rujukan literatur terdahulu. Penggunaan algoritma Naïve Bayes dan *Support Vector Machine* (SVM) terbukti memiliki ketangguhan kompetitif dalam memetakan sentimen publik pada domain olahraga populer [6]. Algoritma Naïve Bayes juga tervalidasi mampu menghasilkan tingkat akurasi tinggi pada analisis sentimen performa tim nasional sepak bola [7]. Selain itu, teknik pra-pemrosesan yang diintegrasikan dengan ekstraksi fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) sangat krusial dalam menangani teks yang memuat ragam istilah teknis [8]. Pendekatan penambahan teks secara umum juga terbukti andal untuk membedah polaritas opini pada isu sensitif berskala nasional [9].

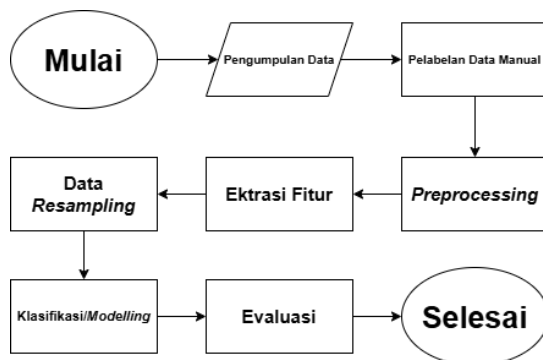
Namun, celah penelitian dari studi-studi tersebut adalah ketiadaan penanganan spesifik terhadap tingginya rasio *out-of-vocabulary* (OOV) pada korpus teks lintas budaya. Komentar warganet dalam kasus Byon Combat memiliki kompleksitas tinggi karena mencampuradukkan bahasa baku, *slang* dua negara, dan jargon teknis olahraga tarung secara bersamaan. Kegagalan komputasi dalam memproses *noise* linguistik tersebut akan mendegradasi akurasi klasifikasi secara fatal. Sebagai rencana pemecahan masalah, penelitian ini mengimplementasikan *custom dictionary-based normalization* sebelum tahap pembobotan fitur guna mereduksi persentase kata yang tidak dikenali mesin secara presisi. Selain itu, potensi ketidakseimbangan kelas (*class imbalance*) yang lazim muncul pada data sentimen akan ditangani menggunakan teknik penambahan proporsi kelas minoritas untuk mencegah bias komputasi [10].

Efektivitas pemrosesan ini dievaluasi dengan membandingkan dua arsitektur klasifikasi utama, yaitu Naïve Bayes yang beroperasi berdasarkan asumsi probabilitas kondisional, dan SVM yang beroperasi menggunakan *linear kernel* untuk mencari batas *hyperplane* optimal. Tujuan akhir penelitian ini adalah mengevaluasi efektivitas penggunaan kamus kustom, mengukur performa kedua algoritma melalui instrumen *confusion matrix*, dan memetakan kecenderungan polaritas sentimen publik terhadap kontroversi hasil pertandingan secara akurat.

B. METODE PENELITIAN

Penelitian ini mengadopsi pendekatan *text mining* berbasis supervised *machine learning* [5]. Ruang lingkup observasi difokuskan pada penggalian data tekstual dari tiga platform media sosial, yaitu

Instagram, TikTok, serta YouTube, yang memuat opini khalayak terkait kontroversi pertandingan Byon Combat Showbiz volume enam. Pengambilan data dieksekusi secara independen secara daring dengan lokasi komputasi di wilayah Bandung. Secara keseluruhan, alur kerja sistematis yang diterapkan dalam studi ini divisualisasikan pada Gambar 1.



Gambar 1. Kerangka Penelitian

Perangkat lunak operasional yang digunakan mencakup ekstensi peramban *Instant Data Scraper* untuk ekstraksi korpus. Seluruh arsitektur pembelajaran mesin dibangun menggunakan bahasa pemrograman Python melalui editor Visual Studio Code. Ekosistem Python terbukti andal dalam mengeksekusi algoritma *machine learning* untuk kebutuhan pemrosesan pemodelan yang kompleks [11]. Pustaka utama yang diimplementasikan meliputi *Scikit-learn* untuk pembobotan dan algoritma, *Sastrawi* untuk manipulasi morfologi leksikal, serta kamus kustom (*custom dictionary*) yang dirangkai secara mandiri.

Teknik pengumpulan data primer dilakukan melalui metode *web scraping*. Proses filtrasi dan *feature selection* diterapkan untuk membuang atribut metadata teknis, sehingga menyisakan 1.081 dokumen komentar unik. Data

tersebut kemudian dilabeli secara manual ke dalam tiga kelas polaritas (positif, negatif, dan netral) oleh dua penilai independen (*annotator*). Keandalan hasil pelabelan ini diukur menggunakan koefisien validasi *Cohen's kappa* [12] untuk memastikan kelayakannya sebagai acuan kebenaran (*ground truth*).

Sebelum dikonversi menjadi representasi numerik, data mentah dibersihkan melalui enam tahapan *preprocessing*. Fase *cleansing* mereduksi karakter *non-alfabetis*, yang dilanjutkan dengan *case folding* guna menyeragamkan teks menjadi huruf kecil. Kalimat lalu dipecah menjadi unit kata mandiri melalui tahap *tokenizing*. Untuk menangani *noise* berupa *slang* dan dialek prokem lintas negara, diterapkan *word normalization* berbasis kamus kustom. Tahap ini krusial untuk menurunkan rasio kata tidak teridentifikasi (*out-of-vocabulary*). Pembersihan diakhiri dengan *stopword removal* untuk mengeliminasi konjungsi tak bermakna, serta *stemming* untuk mereduksi kata berimbuhan menjadi lema dasar.

Dokumen bersih ditransformasikan menjadi matriks bobot menggunakan metode TF-IDF. Guna mencegah terjadinya fenomena kebocoran data (*data leakage*), dataset dipartisi sejak awal menjadi 80% data latih dan 20% data uji. Mengingat distribusi kelas sangat timpang, metode *Synthetic Minority Over-sampling Technique* (SMOTE) [10] diaplikasikan secara eksklusif pada kompartemen data latih. Teknik *resampling* ini membangkitkan data buatan (*synthetic data*) secara geometris guna menyeimbangkan proporsi kelas tanpa memicu *overfitting*.

Teknik analisis data dieksekusi dengan membandingkan dua algoritma klasifikasi. Algoritma pertama adalah Naïve Bayes tipe multinomial yang mengalkulasi probabilitas teoretis fitur, dan algoritma kedua adalah SVM dengan

linear kernel yang bekerja memetakan *hyperplane* pembatas maksimal pada ruang berdimensi tinggi. Evaluasi kinerja dari kedua model klasifikasi tersebut diukur menggunakan matriks kontingensi (*confusion matrix*) untuk menghasilkan nilai akurasi, presisi, *recall*, serta F1-score.

C. HASIL DAN PEMBAHASAN

Pengumpulan Data dan Validasi Pelabelan

Tahap awal mengekstraksi opini publik dari YouTube, TikTok, dan Instagram menghasilkan *dataset* final sebanyak 1.081 dokumen komentar unik pasca-pembersihan *noise* teknis. Data ini kemudian didistribusikan ke dalam tiga polaritas sentimen oleh dua anotator independen sebagai acuan kebenaran (*ground truth*). Matriks kesepakatan pelabelan disajikan pada Tabel 1.

Tabel 1. Matriks Kontingensi Kesepakatan Pelabelan

		Anotator 1			Total
		Negatif	Netral	Positif	
Anotator 2	Negatif	677	28	19	724
	Netral	20	77	19	116
	Positif	0	36	205	241
Total		697	141	243	1.081

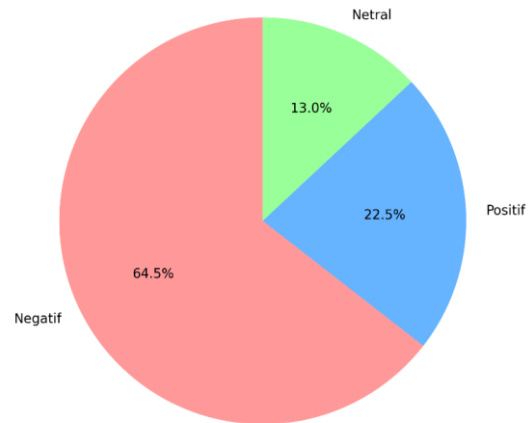
Reliabilitas antar-anotator diuji menggunakan persamaan *Cohen's Kappa* (κ):

$$\kappa = \frac{p_o - p_e}{1 - p_e} \quad (1)$$

Kalkulasi probabilitas persetujuan observasi (p_o) dan persetujuan teoretis (p_e) dari Tabel 1 menghasilkan nilai koefisien $\kappa = 0,776$. Nilai ini melampaui

ambang batas 0,60, mengonfirmasi tingkat kesepakatan substansial (*substantial agreement*).

Visualisasi polaritas opini pada Gambar 2 menunjukkan ketidakseimbangan kelas (*class imbalance*) yang sangat ekstrem.



Gambar 2. Distribusi Sentimen

Publik

Sentimen negatif mendominasi opini warganet sebesar 64,5% (697 data), disusul sentimen positif 22,5% (243 data), dan netral 13,0% (141 data).

Pra-pemrosesan Data dan Ekstraksi Fitur

Enam tahapan linguistik komputasional dieksekusi secara runtut: *cleansing*, *case folding*, *tokenizing*, normalisasi, *stopword removal*, dan *stemming*. Tahapan ini mengimplementasikan *Custom Dictionary-based Normalization* untuk memetakan dan mereduksi kata *Out Of Vocabulary* (OOV) berupa *slang* dan jargon olahraga kombat ke dalam bentuk baku.

Tabel 2. Perbandingan Teks Asli dan Hasil Akhir Pra-pemrosesan

Teks Asli	Hasil Akhir Pra-pemrosesan
ini serius PRINCE	serius putra imbang

DRAW?	
"Di ring cuma kau dan aku je" nyatanya 1 keluarga ikut'an juga wak🤔😊	ring cuma kau saya nyata keluarga ikut an kawan

Transformasi ke dalam ruang vektor numerik dioperasikan menggunakan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF). Guna mengeliminasi *data leakage*, 20% data (217 dokumen) diisolasi secara mutlak sebagai kompartemen data uji. Ekstraksi matriks kosakata sepenuhnya dibentuk berdasarkan probabilitas dari 80% data latih (864 dokumen).

Tabel 3. Sampel Distribusi Bobot TF-IDF

Kata	Bobot TF-IDF
serius	0,798925
imbang	0,497142
putra	0,338481

Nilai bobot pada Tabel 3 mengonfirmasi kapabilitas TF-IDF dalam menonjolkan fitur kata bermakna (seperti "serius" dan "imbang") berdasarkan tingkat kelangkaannya secara global (IDF)

Penanganan Ketidakseimbangan Kelas (Resampling)

Distribusi 864 dokumen latih yang dikuasai kelas negatif (544 data) berisiko memicu *overfitting*. Teknik *Synthetic Minority Over-sampling Technique* (SMOTE) diterapkan pada data latih guna menyamaratakan distribusi kelas minoritas. SMOTE melakukan interpolasi geometris secara cerdas sehingga probabilitas dari seluruh kelas sentimen terdistribusi ekuivalen masing-masing menjadi 544 dokumen tanpa menduplikasi nilai aktual.

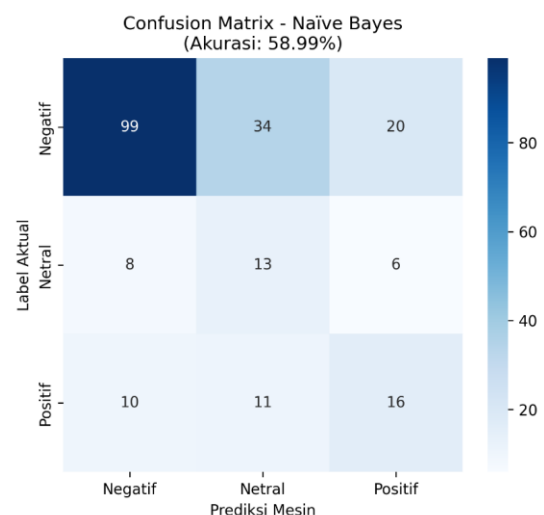
Tabel 4. Distribusi Kelas Data Latih Sebelum dan Sesudah SMOTE

Kelas Sentimen	Sebelum SMOTE	Sesudah SMOTE
Negatif	544	544
Positif	206	544
Netral	114	544
Total	864	1.632

Berdasarkan Tabel 2, penerapan SMOTE berhasil menyamaratakan distribusi seluruh kelas menjadi 544 dokumen per kategori. Pemerataan ini memberikan landasan ruang probabilitas yang seimbang bagi algoritma dalam mengklasifikasikan sentimen.

Hasil Klasifikasi dan Evaluasi Model

Kinerja algoritma Naïve Bayes tipe multinomial dan *Support Vector Machine* (SVM) berbasis *Linear Kernel* dievaluasi murni terhadap 217 dokumen data uji steril. Perbandingan misklasifikasi dari kedua mesin divisualisasikan pada Gambar 3 dan Gambar 4.



Gambar 3. *Confusion Matrix Naïve Bayes*

merupakan bentuk protes teknis terhadap kualitas regulasi dan keputusan pengadilan, yang direpresentasikan secara absolut oleh tingginya frekuensi kata "aturan", "imbang", "pelanggaran", "poin", dan "wasit". Implementasi *Custom Dictionary-based Normalization* pada tahap pra-pemrosesan terbukti sangat krusial dalam menormalisasi *noise* linguistik (*slang* siber) dan mereduksi rasio kata *Out-Of-Vocabulary* (OOV) pada teks *sportainment*. Pada fase pemodelan dengan data uji steril, algoritma *Support Vector Machine* (SVM) tipe *Linear Kernel* tervalidasi sebagai model paling optimal dengan tingkat akurasi 59,45%, mengungguli Naïve Bayes (58,99%) dalam membedah matriks berdimensi tinggi yang ruang probabilitas kelasnya telah diseimbangkan menggunakan teknik SMOTE.

Saran

Berdasarkan evaluasi terhadap limitasi model, terdapat beberapa rekomendasi optimasi komputasi yang ringan untuk pengembangan penelitian analitik sentimen selanjutnya:

1. Ekspansi N-Gram pada Ekstraksi Fitur: Stagnasi akurasi menunjukkan metode unigram tidak cukup untuk menangkap sarkasme. Penelitian lanjutan direkomendasikan untuk mengintegrasikan kombinasi bigram atau trigram pada pembobotan TF-IDF guna mempertahankan struktur kontekstual frasa tanpa membebani memori perangkat secara signifikan.
2. Optimasi Hyperparameter: Untuk meningkatkan performa identifikasi, disarankan melakukan optimasi hyperparameter (seperti penerapan GridSearchCV) guna mencari parameter C dan gamma terbaik pada arsitektur SVM, atau melakukan komparasi tambahan dengan algoritma

ansambel yang hemat komputasi seperti Random Forest.

3. Komparasi Lintas *Platform*: Perluasan data observasi secara paralel menggunakan dataset dari platform dengan karakteristik distribusi informasi berbeda (seperti X/Twitter) diperlukan untuk memetakan disparitas gaya bahasa antar-ekosistem digital.
4. Pemutakhiran Kamus Kustom: Pengembangan korpus kamus kustom secara dinamis direkomendasikan untuk terus mengadopsi evolusi serapan istilah olahraga tarung internasional serta modifikasi dialek slang lintas negara.

E. REFERENSI

- [1] K. Sikumbang dkk., "Peranan Media Sosial Instagram terhadap Interaksi Sosial dan Etika pada Generasi Z," *joe*, vol. 6, no. 2, hlm. 11029–11037, Jan 2024, doi: 10.31004/joe.v6i2.4888.
- [2] A. Richelieu, "From sport to 'sportainment': The art of creating an added-value brand experience for fans," *Journal of Brand Strategy*, vol. 9, no. 4, hlm. 408–422, Jan 2021.
- [3] Liputan6.com, "Byon Combat Series 6 Penuh Kontroversi, Indonesia Telan Kekalahan dari Malaysia," liputan6.com. Diakses: 19 Desember 2025. [Daring]. Tersedia pada: <https://www.liputan6.com/bola/read/6218635/byon-combat-series-6-penuh-kontroversi-indonesia-telan-kekalahan-dari-malaysia>
- [4] R. K. Permana, M. A. Senubekti, N. Anggraini, dan M. S. Usman, "PREDIKSI KELULUSAN MAHASISWA DENGAN MENGGUNAKAN ALGORITMA DECISION TREE," *JATI (Jurnal*

- Mahasiswa Teknik Informatika*), vol. 9, no. 5, hlm. 7744–7751, Jul 2025, doi: 10.36040/jati.v9i5.14877.
- [5] W. Fan, L. Wallace, S. Rich, dan Z. Zhang, “Tapping the power of text mining,” *Commun. ACM*, vol. 49, no. 9, hlm. 76–82, Sep 2006, doi: 10.1145/1151030.1151032.
- [6] S. Fathurrohman, I. R. Afandi, dan F. N. Hasan, “Sentiment Analysis of TIMNAS Indonesia’s Participation in the Asian Cup U23 2024 on X Using Naïve Bayes and SVM,” *IJID (International Journal on Informatics for Development)*, vol. 13, no. 1, hlm. 434–447, Agu 2024, doi: 10.14421/ijid.2024.4504.
- [7] W. Alfiyani, D. A. Fatah, dan F. Irhamni, “Penerapan algoritma Naïve Bayes untuk analisis sentimen pada media sosial X terhadap performa tim nasional sepak bola Indonesia di era kepemimpinan Shin Tae-yong,” *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 3, hlm. 3969–3977, 2025.
- [8] R. N. Rahman, A. Rahim, dan W. J. Pranoto, “Analisis sentimen ulasan game efootball 2024 pada playstore menggunakan algoritma naïve bayes,” *JURNAL ILMIAH INFORMATIKA*, vol. 13, no. 01, hlm. 38–44, Mar 2025, doi: 10.33884/jif.v13i01.9913.
- [9] N. Anggraini, E. S. N. Harahap, dan T. B. Kurniawan, “Text Mining - Analisis Teks Terkait Isu Vaksinasi COVID-19 (Text Mining - Text Analysis Related to COVID-19 Vaccination Issues),” *IPTEKKOM*, vol. 23, no. 2, hlm. 141–153, Des 2021, doi: 10.17933/iptekkom.23.2.2021.141-153.
- [10] N. V. Chawla, K. W. Bowyer, L. O. Hall, dan W. P. Kegelmeyer, “SMOTE: Synthetic Minority Over-sampling Technique,” *Journal of Artificial Intelligence Research*, vol. 16, hlm. 321–357, Jun 2002, doi: 10.1613/jair.953.
- [11] N. Febriana, F. Fadzira, M. A. Senubekti, dan R. Suharsih, “Prediksi Penilaian Kinerja Hakim Dengan Penerapan Machine Learning Menggunakan Tools Python,” *TEKNO*, vol. 2, no. 1, hlm. 44–51, Mar 2024, doi: 10.62565/tekno.v2i1.23.
- [12] J. Cohen, “A Coefficient of Agreement for Nominal Scales,” *Educational and Psychological Measurement*, vol. 20, no. 1, hlm. 37–46, Apr 1960, doi: 10.1177/001316446002000104.