

# Implementasi Algoritma Naïve Bayes untuk Analisis Sentimen Komentar Netizen X Mengenai Pemecatan Patrick Kluivert dari Timnas Indonesia

I Made Dwi Septayuda<sup>1</sup>, Dolly Virgian Shaka Yudha Sakti<sup>2\*</sup>, Nofiani<sup>3</sup>

<sup>1,2,3</sup> Fakultas Teknologi Informasi, Teknik Informatika, Universitas Budi Luhur, DKI Jakarta, Indonesia  
Jl. Ciledug Raya, Petungkang Utara, Jakarta Selatan, 12260. DKI Jakarta  
E-mail: <sup>1</sup> 2111500936@student.budiluhur.ac.id, <sup>2\*</sup> dolly.virgianshaka@budiluhur.ac.id, <sup>3</sup> nofiani@budiluhur.ac.id  
(\*: corresponding author)

**Abstrak—** Penelitian ini mengimplementasikan algoritma Multinomial Naïve Bayes untuk memetakan sentimen netizen di media sosial X mengenai pemberhentian Patrick Kluivert dari kursi pelatih Timnas Indonesia. Menggunakan korpus sebanyak 351 cuitan yang dijangkau dengan kata kunci "Kluivert Out", data diproses melalui tahapan text mining yang komprehensif, mulai dari preprocessing hingga ekstraksi fitur TF-IDF. Hasil pengujian dengan rasio pembagian data 90:10 mengungkapkan bahwa mayoritas opini bersifat negatif (59,54%), yang merefleksikan rasa kecewa publik terhadap keputusan federasi. Model ini mencapai tingkat akurasi 83,81% dengan performa sangat optimal pada deteksi kelas negatif melalui nilai recall 95%. Namun, faktor data imbalance menjadi kendala utama dalam mengklasifikasikan kategori netral secara akurat. Studi ini membuktikan bahwa metode Naïve Bayes fungsional dalam menangkap arus utama aspirasi masyarakat bola Indonesia di era transformasi digital.

**Kata Kunci—** Multinomial Naïve Bayes, Analisis Sentimen, TF-IDF, Media Sosial X, Patrick Kluivert

**Abstract—** This study implements the Multinomial Naïve Bayes algorithm to map netizen sentiment on social media platform X regarding the dismissal of Patrick Kluivert from his position as the head coach of the Indonesian National Team. Utilizing a corpus of 351 tweets collected via the keyword "Kluivert Out," the data underwent comprehensive text mining stages, from preprocessing to TF-IDF feature extraction. Experimental results using a 90:10 data split ratio reveal that the majority of public opinion is negative (59.54%), reflecting widespread public disappointment with the federation's decision. The classification model achieved an accuracy rate of 83.81%, demonstrating highly optimal performance in detecting the negative class with a recall value of 95%. However, data imbalance remains a significant obstacle in accurately classifying the neutral category. This study proves that the Naïve Bayes method is functional in capturing the mainstream aspirations of the Indonesian football community in the era of digital transformation.

**Keywords—** Multinomial Naïve Bayes, Sentiment Analysis, TF-IDF, Social Media X, Patrick Kluivert

## I. PENDAHULUAN

Media sosial X (Twitter) kini menjadi ruang diskusi favorit masyarakat Indonesia, tak terkecuali soal sepak bola nasional. Ketika Patrick Kluivert dipecat dari posisi pelatih Timnas

Indonesia, platform ini langsung dibanjiri berbagai respons—ada yang mendukung keputusan tersebut, ada pula yang mengkritik habis-habisan. Ribuan komentar bermunculan, menciptakan kumpulan data opini publik yang sangat besar namun mustahil untuk dibaca dan dipahami satu per satu secara manual.

Dengan penetrasi pengguna mencapai 62% dari total populasi Indonesia, platform X memiliki pengaruh yang sangat besar dalam membentuk gerakan opini digital yang mampu memengaruhi kebijakan institusional maupun isu nasional [1], [2]. Karakteristik X yang bersifat terbuka dan menyebarkan informasi secara real-time menjadikannya sumber data berharga yang dapat digunakan untuk memetakan pendapat publik terhadap berbagai macam fenomena sosial dan olahraga [3], [4].

Sepak bola bagi masyarakat Indonesia bukan sekadar olahraga, melainkan fenomena sosial yang memicu reaksi emosional kolektif yang kuat [5], [6]. Keputusan federasi sepak bola Indonesia untuk memberhentikan Patrick Kluivert dari jabatan pelatih Timnas Indonesia seketika memicu banjir respons dari netizen di platform X. Ribuan cuitan bermunculan dalam rentang waktu singkat, mencerminkan berbagai spektrum emosi mulai dari dukungan terhadap keputusan tersebut hingga kritik tajam dan kekecewaan. Bagi pihak berkepentingan seperti PSSI, memahami esensi dari percakapan publik ini sangat krusial sebagai bahan evaluasi kebijakan strategis di masa depan [3]. Tetapi karena jumlah data yang sangat banyak dan tidak terstruktur membuat analisis tanpa teknologi komputer membuat hal ini menjadi tidak efisien dan rentan terhadap subjektivitas.

Untuk mengatasi tantangan analisis manual, teknologi Analisis Sentimen atau opinion mining hadir sebagai solusi ilmiah. Analisis sentimen adalah sebuah penerapan pemrosesan bahasa alami (*Natural Language Processing*) dan analisis teks untuk mengidentifikasi, mengekstrak, dan menganalisis secara sistematis informasi subjektif serta kondisi emosional dalam suatu teks [2]. Teknologi ini memungkinkan pembagian opini / pendapat publik ke dalam beberapa kategori diantaranya sentiment positif, sentiment negatif, atau netral secara otomatis dan objektif [7]. Pemanfaatan analisis sentimen telah terbukti efektif dalam berbagai kasus, mulai dari pemantauan merek,

riset pasar, hingga evaluasi kebijakan publik seperti pembangunan IKN dan Undang-Undang Cipta Kerja [2][3].

Algoritma Naïve Bayes sudah banyak digunakan untuk analisis sentimen di media social, dan telah banyak dieksplorasi dalam berbagai penelitian terdahulu [8]–[10]. Dalam domain sosial-politik, analisis terhadap tagar #kaburajadulu menunjukkan bahwa implementasi Naïve Bayes menghasilkan sentimen negatif publik dengan nilai akurasi 75% [1]. Sementara itu, penelitian mengenai kebijakan pembangunan IKN mencapai akurasi sebesar 61% dalam memetakan dukungan masyarakat [3]. Dalam konteks olahraga dan industri hiburan, algoritma ini juga menunjukkan performa yang kompetitif. Penelitian lain yang membahas komentar pada media sosial Instagram klub Persija Jakarta membuktikan efektivitas Naïve Bayes dalam menangani interaksi suporter sepak bola[5]. Selain itu, penelitian lain membahas tentang kesuksesan di ajang Thomas Cup 2020 yang diraih Indonesia menunjukkan tingkat akurasi mencapai nilai tinggi yaitu 95% [6]. Pada penelitian lain menyatakan bahwa performa model sangat dipengaruhi oleh kualitas preprocessing dan komposisi pembagian data (data splitting), di mana rasio 80:20 sering kali memberikan hasil yang lebih optimal untuk dataset terbatas [11].

Pemilihan algoritma Multinomial Naïve Bayes didasarkan pada spesialisasi metode tersebut dalam mengolah data tekstual menggunakan perhitungan TF-IDF. Keunggulan utama teknik ini terletak pada kecepatan proses serta kebutuhan sampel pelatihan yang minimal dibandingkan model pembelajaran mendalam (deep learning), sehingga tetap relevan untuk menghasilkan pemetaan sentimen yang presisi dalam konteks bahasa Indonesia [5], [12]. Penggunaan teknik ini diharapkan dapat meminimalisir kesalahan informasi dan memberikan gambaran yang akurat mengenai dinamika opini publik terkait manajemen tim nasional [11].

Meskipun metode Naïve Bayes telah teruji pada berbagai topik umum, penerapannya pada isu spesifik pemecatan pelatih Timnas Indonesia memiliki tantangan tersendiri terkait penggunaan bahasa slang sepak bola dan istilah-istilah khas netizen Indonesia. Oleh karena itu, penelitian ini merumuskan masalah mengenai bagaimana performa algoritma Multinomial Naïve Bayes dalam mengklasifikasikan opini netizen X terkait isu pemecatan Patrick Kluivert.

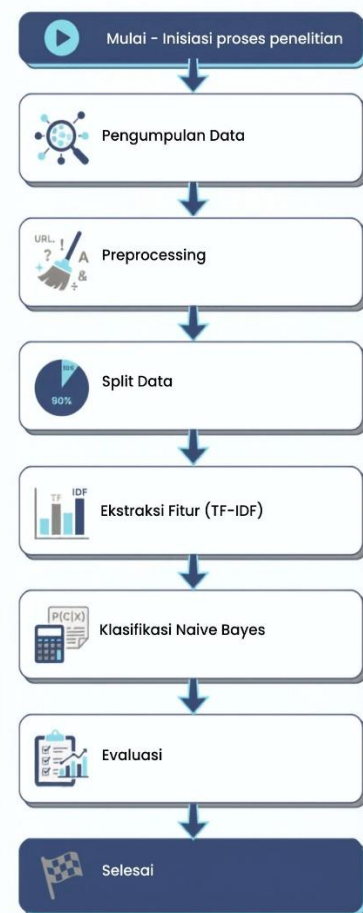
Tujuan utama penelitian ini adalah untuk menyediakan alat analisis yang objektif, efisien, dan terukur dalam memetakan persepsi masyarakat terhadap keputusan strategis manajemen Timnas Indonesia. Hasil penelitian ini diharapkan tidak hanya memberikan kontribusi akademis dalam bidang text mining, tetapi juga menjadi acuan praktis bagi pemangku kebijakan olahraga dalam merumuskan strategi komunikasi publik yang lebih tepat sasaran dan adaptif terhadap opini masyarakat di era digital.

## II. METODE PENELITIAN

Penelitian ini menerapkan metodologi yang sistematis dalam beberapa fase demi menjaga ketepatan pengelompokan sentimen. Fase yang dilalui berawal dari ekstraksi data mentah (*crawling*), pembersihan teks (*preprocessing*), tahapan pelabelan, perhitungan bobot dengan TF-IDF, penerapan

model klasifikasi Multinomial Naïve Bayes, hingga evaluasi performa melalui confusion matrix. Tujuan dari alur ini adalah mengonversi opini bebas yang tidak terstruktur menjadi data kuantitatif yang dapat dinilai secara objektif.

Gambaran dari metode penelitian yang dilakukan pada penelitian ini seperti terlihat pada Gambar 1.



Gambar 1 Metode Penelitian

### A. Pengumpulan Data

Dataset pada penelitian ini dihimpun melalui teknik pengumpulan otomatis pada platform X menggunakan bantuan API resmi dan bahasa pemrograman Python. Masa pengambilan data dibatasi pada periode 11 sampai 20 Oktober 2025, yang bertepatan dengan momentum paling intens terkait isu pemecatan Patrick Kluivert dari kursi pelatih. Pencarian difokuskan pada penggunaan istilah "Kluivert Out" guna memastikan opini yang terkumpul bersifat relevan dengan topik penelitian. Melalui prosedur tersebut, sebanyak 351 cuitan telah disimpan dalam file CSV untuk diproses lebih lanjut melalui sistem yang telah dibangun.

### B. Pemrosesan Awal (Text Preprocessing)

Data mentah yang dihimpun dari media sosial cenderung tidak terstruktur dan sarat akan gangguan (*noise*), maka diperlukan fase pembersihan teks secara intensif agar data tersebut layak diproses oleh algoritma. Rangkaian tahapan preprocessing dalam riset ini diimplementasikan melalui prosedur berikut [13]:

- 1) *Cleansing*: Ekstraksi teks murni dengan membuang komponen tambahan seperti URL, mention, hashtag, angka, serta tanda baca yang tidak diperlukan.
- 2) *Case Folding*: Penyeragaman tipografi teks ke dalam format huruf kecil untuk menghindari ambiguitas interpretasi oleh algoritma.
- 3) *Tokenizing*: Dekomposisi teks dari bentuk kalimat menjadi kumpulan kata atau token yang mandiri.
- 4) *Stopword Removal*: Penyeleksian kata untuk membuang istilah-istilah generik yang tidak memiliki makna signifikan dalam analisis emosi publik.
- 5) *Stemming*: Penerapan algoritma klasifikasi untuk mentransformasi kata berimbuhan kembali ke bentuk kata dasar (akar kata) guna meminimalisir keragaman fitur morfologis.

### C. Split Data

Tahap berikutnya setelah prapemrosesan adalah pemisahan dataset menjadi kategori data latih dan data uji dengan skema pembagian 90:10. Alokasi 90% atau setara dengan 316 cuitan ditujukan untuk melatih kemampuan sistem dalam mengenali karakteristik opini, sedangkan 10% lainnya (35 cuitan) disimpan guna mengevaluasi akurasi prediksi pada data baru. Keputusan untuk memperbesar porsi data pelatihan didasari oleh volume dataset yang relatif minim agar model klasifikasi mendapatkan pengetahuan yang optimal dalam mempelajari kelas-kelas minoritas. Hal ini krusial dilakukan terutama untuk memperkuat deteksi pada sentimen netral dan positif yang distribusinya tidak seimbang dalam dataset.

### D. Ekstraksi Fitur (TF-IDF)

Setelah dataset terbagi, informasi tekstual ditransformasikan menjadi format angka melalui pendekatan *Term Frequency-Inverse Document Frequency* (TF-IDF). Teknik ini menetapkan signifikansi tiap kata dengan mengukur intensitas kemunculannya di tingkat dokumen (*Term Frequency*) relatif terhadap tingkat kelangkaannya di korpus data secara keseluruhan (*Inverse Document Frequency*) [14]–[16]. Fitur-fitur yang memiliki skor TF-IDF unggul diidentifikasi sebagai elemen kunci yang mengandung muatan emosional atau sentimen dalam cuitan yang dianalisis.

### E. Klasifikasi Naïve Bayes

Penelitian ini memanfaatkan Multinomial Naïve Bayes sebagai mesin klasifikasi utama karena efektivitasnya dalam menangani data tekstual berbasis bobot kemunculan istilah. Prosedur penghitungan dimulai dengan menentukan probabilitas prior untuk setiap kategori sentimen dan dilanjutkan dengan kalkulasi probabilitas *likelihood* untuk setiap fitur kunci. Untuk menjamin stabilitas perhitungan numerik, diterapkan skema *Laplace Smoothing* yang berfungsi mencegah kegagalan klasifikasi akibat probabilitas nol pada kata yang absen [17]–[19]. Keputusan sistem dalam melabeli opini ke dalam kelompok Positif, Negatif, atau Netral sepenuhnya merujuk pada nilai probabilitas posterior tertinggi yang dihasilkan oleh model.

### F. Evaluasi

Tahap terakhir adalah mengevaluasi performa model klasifikasi menggunakan *Confusion Matrix*. Melalui alat evaluasi ini, akan dihitung empat metrik penilaian objektif untuk mengukur efektivitas sistem:

- 1) *Akurasi*: Mengukur tingkat ketepatan prediksi model secara keseluruhan dibandingkan dengan total data.
- 2) *Presisi*: Menunjukkan tingkat keandalan model dalam memprediksi setiap kelas sentimen secara tepat.
- 3) *Recall*: Mengukur kemampuan model dalam mengidentifikasi kembali seluruh data yang termasuk dalam kelas tertentu.
- 4) *F1-Score*: Nilai rata-rata harmonik antara presisi dan *recall* yang memberikan gambaran keseimbangan performa model, terutama pada dataset dengan distribusi kelas yang tidak seimbang.

## III. HASIL DAN PEMBAHASAN

### A. Lingkungan Implementasi Sistem

Keberhasilan implementasi model analisis sentimen sangat bergantung pada infrastruktur teknis yang digunakan. Lingkungan implementasi dalam penelitian ini mencakup spesifikasi perangkat lunak (software) dan perangkat keras (hardware) yang dipilih untuk memastikan proses pengolahan data teks dan eksekusi algoritma Multinomial Naïve Bayes berjalan secara efisien dan optimal.

#### 1) Spesifikasi Perangkat Lunak (*Software*)

Perangkat lunak berperan sebagai instruksi logis untuk menjalankan algoritma dan memproses dataset. Spesifikasi perangkat lunak yang digunakan dalam pengembangan aplikasi ini meliputi:

**Sistem Operasi:** Windows 10, dipilih sebagai platform dasar karena stabilitasnya dalam menjalankan berbagai aplikasi pengembang.

**Bahasa Pemrograman:** Python, digunakan sebagai bahasa utama karena memiliki ekosistem pustaka (library) yang sangat kaya untuk kebutuhan *Natural Language Processing* (NLP) dan *machine learning*.

**Integrated Development Environment (IDE):** Visual Studio Code, dimanfaatkan sebagai lingkungan pengembangan terpadu untuk menulis, melakukan debugging, dan mengelola kode program.

**Browser:** Google Chrome, digunakan untuk mengakses layanan komputasi awan seperti Google Colab saat melakukan proses crawling data dari platform X serta untuk pengujian antarmuka aplikasi.

#### 2) Spesifikasi Perangkat Keras (*Hardware*)

Dukungan perangkat keras yang memadai sangat krusial, terutama saat melakukan transformasi data teks (TF-IDF) dan proses pelatihan model yang membutuhkan memori stabil. Perangkat keras yang digunakan memiliki spesifikasi sebagai berikut:

**Prosesor:** Intel Core i5-10500H, yang memberikan daya komputasi yang cukup kuat untuk menangani proses stemming dan klasifikasi pada dataset.

Memori (RAM): 8 GB, untuk memastikan kelancaran proses multitasking saat aplikasi dan IDE berjalan secara bersamaan.

Penyimpanan (Storage): SSD 477 GB, yang memberikan kecepatan akses data yang tinggi saat pembacaan dan penulisan file dataset dalam format CSV.

Kartu Grafis: NVIDIA Geforce GTX 1650, yang mendukung performa sistem secara keseluruhan selama proses pengembangan berlangsung.

## B. Analisis Karakteristik Dataset dan Distribusi Sentimen

Tahap analisis karakteristik dataset ini bertujuan untuk memberikan gambaran mengenai data yang telah berhasil dikumpulkan serta bagaimana persepsi netizen X terhadap isu pemecatan Patrick Kluivert.

### 1) Karakteristik Dataset Hasil Crawling

Data penelitian ini merupakan data primer yang diperoleh secara langsung dari platform X (Twitter) menggunakan teknik crawling melalui layanan Google Colab. Parameter pencarian difokuskan pada kata kunci "Kluivert Out" untuk menangkap reaksi publik yang spesifik terhadap isu tersebut. Data dikumpulkan selama periode 11 hingga 20 Oktober 2025, yang merupakan masa puncak viralnya diskusi mengenai kebijakan manajemen tim nasional. Dari proses tersebut, berhasil terkumpul sebanyak 351 cuitan yang relevan untuk dianalisis lebih lanjut.

### 2) Pelabelan Data Manual

Sebelum dilakukan pengolahan oleh algoritma, seluruh komentar dalam dataset tersebut melalui proses pelabelan manual untuk menjamin akurasi dan konsistensi hasil klasifikasi. Cuitan dikelompokkan ke dalam tiga kategori utama berdasarkan muatan emosional dan isi pesan yang disampaikan:

**Sentimen Positif:** Mencakup komentar yang berisi dukungan, apresiasi, atau harapan optimis terhadap keputusan pemecatan Patrick Kluivert.

**Sentimen Netral:** Meliputi komentar yang bersifat informatif, objektif, berupa pertanyaan, atau penyampaian fakta tanpa menunjukkan emosi mendukung maupun menolak secara tegas.

**Sentimen Negatif:** Berisi kritik, rasa frustrasi, kekecewaan, atau penolakan terhadap keputusan yang diambil oleh federasi.

### 3) Distribusi Sentimen Publik

Berdasarkan hasil pelabelan manual terhadap 351 data tersebut, ditemukan kecenderungan opini publik yang sangat kontras. Distribusi frekuensi sentimen dapat dilihat pada Tabel 1.

TABEL 1  
DISTRIBUSI DATA KESELURUHAN

| Kategori | Jumlah Data | Presentase | Karakteristik         |
|----------|-------------|------------|-----------------------|
| Positif  | 107         | 30.48%     | Mendukung keputusan   |
| Netral   | 45          | 12.82%     | Objektif / Informatif |
| Negatif  | 209         | 59.54%     | Menolak keputusan     |
| Total    | 351         | 100%       | -                     |

Hasil analisis menunjukkan bahwa sentimen negatif mendominasi percakapan publik dengan persentase mencapai 59,54%. Dominasi ini mengindikasikan adanya rasa kekecewaan dan penolakan yang kuat dari mayoritas netizen terhadap keputusan pemecatan tersebut. Di sisi lain, sentimen positif tercatat sebesar 30,48% dan sentimen netral menjadi kategori dengan jumlah paling sedikit, yakni hanya 12,82%. Tingginya volume opini negatif ini mencerminkan dinamika emosi kolektif masyarakat sepak bola Indonesia yang aktif menyuarakan kritik melalui platform digital.

## C. Analisis Hasil Preprocessing dan Pembobotan Kata

Tahap ini merupakan fase krusial dalam mengubah data tidak terstruktur (unstructured data) dari media sosial X menjadi representasi numerik yang dapat dipahami oleh algoritma Multinomial Naïve Bayes. Proses ini terdiri dari pembersihan teks secara bertahap dan perhitungan bobot fitur menggunakan metode TF-IDF.

### 1) Analisis Transformasi Teks (Preprocessing)

Data mentah yang diperoleh dari hasil crawling mengandung banyak gangguan (noise) seperti simbol, URL, dan penulisan yang tidak seragam. Melalui lima tahap preprocessing, teks dibersihkan agar hanya menyisakan kata-kata yang bermakna sentimen:

**Cleansing:** Sistem berhasil menghapus elemen non-teks. Sebagai contoh, cuitan mentah "@PSSI Kluivert out ->" PHS in"" dibersihkan dari simbol mention akun menjadi "Kluivert out ->" PHS in"".

**Case Folding:** Semua karakter diubah menjadi huruf kecil untuk menjaga konsistensi. Kata "Kluivert" dan "PHS" diseragamkan menjadi "kluivert" dan "phs".

**Tokenizing:** Kalimat dipecah menjadi unit kata tunggal untuk mempermudah penghitungan frekuensi.

**Stopword Removal:** Kata-kata umum yang sering muncul namun tidak memiliki muatan informasi (seperti "di", "dan", "itu") dihapus.

**Stemming:** Kata berimbuhan dikembalikan ke bentuk dasarnya. Misalnya, dalam dataset ditemukan kata "dipecat" diubah menjadi kata dasar "pecat" dan "balikin" menjadi "balik".

Hasil akhir dari tahap ini adalah teks yang bersih dan siap dihitung bobotnya. Perbandingan teks asli dan teks hasil preprocessing dapat dilihat pada antarmuka aplikasi yang telah dibangun.

### 2) Analisis Ekstraksi Fitur dan Pembobotan Kata (TF-IDF)

Setelah melewati fase pembersihan, data tekstual ditransformasikan ke dalam format vektor numerik melalui metode Term Frequency-Inverse Document Frequency (TF-IDF). Pendekatan ini berfungsi untuk mengukur tingkat kepentingan sebuah istilah dengan membandingkan intensitas kemunculannya dalam satu cuitan terhadap prevalensi kata tersebut di seluruh korpus data. Dalam operasionalnya, sistem melakukan kalkulasi otomatis sebagai berikut:

**Term Frequency (TF):** Menakar kekerapan kata kunci di tingkat dokumen individu. Sebagai ilustrasi, pada satu sampel uji, kata "Kluivert" tercatat memiliki nilai TF sebesar 0,0333.

*Inverse Document Frequency* (IDF): Mengevaluasi tingkat kelangkaan fitur kata di dalam dataset global, di mana bobot yang lebih tinggi diberikan pada istilah yang jarang muncul. Merujuk pada data penelitian, kata "Patrick" mendapatkan skor IDF 1,3681, sedangkan "Kluivert" berada di angka 1,0488.

Skor TF-IDF: Merupakan produk perkalian antara nilai TF dan IDF yang menentukan signifikansi akhir sebuah fitur. Berdasarkan hasil pemrosesan model, bobot akhir untuk kata "Kluivert" adalah 0,0456 dan untuk kata "Patrick" mencapai 0,0350.

### 3) Visualisasi Kata Kunci Dominan

Melalui algoritma pembobotan tersebut, sistem mengidentifikasi sejumlah kata kunci utama yang merepresentasikan opini netizen terhadap isu ini. Berdasarkan tampilan Top 50 Vocabulary pada aplikasi, kata-kata yang paling sering muncul dan memiliki bobot signifikan meliputi: "kluivert", "out", "patrick", "pecat", "latih", "timnas", "indonesia", "pssi", "erick", dan "sty".

Dominasi kata-kata seperti "out", "pecat", dan "kluivert" menunjukkan secara visual bahwa topik pembicaraan netizen memang terpusat pada desakan atau reaksi negatif terkait posisi kepelatihan Patrick Kluivert, yang selaras dengan temuan distribusi sentimen negatif yang mencapai 59,5%.

### D. Implementasi Model Klasifikasi Naïve Bayes

Tahap implementasi klasifikasi merupakan inti dari sistem analisis sentimen ini, di mana algoritma Multinomial Naïve Bayes digunakan untuk mempelajari pola dari dataset yang telah dibobotkan. Mengingat jumlah dataset yang relatif terbatas (351 cuitan), penelitian ini secara strategis menggunakan porsi 90% data untuk pelatihan (315 cuitan) dan 10% untuk pengujian. Pilihan ini diambil untuk memberikan peluang lebih besar bagi model dalam mengenali karakteristik kelas minoritas, seperti sentimen netral dan positif, yang seringkali memiliki jumlah data jauh lebih sedikit dibandingkan kelas negatif.

#### 1) Perhitungan Probabilitas Prior (*Prior Probability*)

Langkah pertama dalam pembangunan model adalah menghitung *Prior Probability*, yaitu kemungkinan awal kemunculan setiap kelas sentimen dalam dataset pelatihan. Nilai ini diperoleh dengan membagi jumlah data pada masing-masing kelas dengan total data latih. Berdasarkan distribusi data pada 315 data latih, diperoleh hasil sebagai berikut:

- $P(\text{Negatif}) = 209 / 315 = 0,663$
- $P(\text{Positif}) = 107 / 315 = 0,339$
- $P(\text{Netral}) = 45 / 315 = 0,142$

#### 2) Pembentukan Model Likelihood dengan Laplace Smoothing

Tahap selanjutnya adalah menghitung Likelihood, yaitu probabilitas munculnya sebuah kata kunci pada kelas tertentu. Untuk menghindari nilai probabilitas nol pada kata yang tidak ditemukan dalam salah satu kelas, sistem menerapkan teknik Laplace Smoothing. Dalam perhitungan ini, sistem mengidentifikasi jumlah kata unik (vocabulary size) sebanyak 1.317 kata. Total kata pada setiap kelas diidentifikasi sebagai dasar pembagi: kelas negatif (1.225 kata), positif (778 kata),

dan netral (338 kata). Contoh hasil perhitungan probabilitas beberapa kata kunci dominan pada kelas negatif meliputi:

$$P(\text{Pecat} | \text{Negatif}) = (160+1) / (1.225 + 1.317) = 0,0633$$

$$P(\text{Kluivert} | \text{Negatif}) = (150+1) / (1.225 + 1.317) = 0,0594$$

$$P(\text{Out} | \text{Negatif}) = (140+1) / (1.225 + 1.317) = 0,0554$$

#### 3) Tahap Klasifikasi Komentar (*Posterior Probability*)

Proses klasifikasi akhir dilakukan dengan menjumlahkan nilai log dari Prior Probability dan log Likelihood dari seluruh kata yang muncul dalam satu cuitan untuk mendapatkan nilai Posterior Probability. Sistem akan memilih label kelas yang memiliki nilai skor tertinggi. Sebagai ilustrasi, sistem melakukan pengujian pada satu cuitan sampel. Berdasarkan hasil komputasi logaritma, didapatkan skor perbandingan sebagai berikut:

$$\text{Total skor log-negatif} = -30.885$$

$$\text{Total skor log-positif} = -33.984$$

$$\text{Total skor log-netral} = -40.674$$

Hasil tersebut menunjukkan bahwa nilai terbesar berada pada kelas NEGATIF (-30.885), sehingga sistem secara otomatis mengklasifikasikan cuitan tersebut ke dalam kategori sentimen negatif. Pendekatan berbasis logaritma ini digunakan untuk menjaga stabilitas perhitungan numerik dan mencegah kesalahan pembulatan pada angka probabilitas yang sangat kecil.

### E. Evaluasi Performa Model

Evaluasi performa merupakan tahap krusial untuk mengukur sejauh mana model Multinomial Naïve Bayes mampu menggeneralisasi pola yang telah dipelajari dari data latih ke data baru. Data uji yang digunakan pada pengujian ini sebanyak 10% dari dataset yaitu 36 cuitan.

#### 1) *Confusion Matrix*

Berdasarkan pengujian dengan rasio 90:10 menggunakan program yang telah dibuat, diperoleh distribusi prediksi seperti terlihat pada Gambar 2.

**Confusion Matrix**

| Actual / Predicted | Pred: negatif | Pred: netral | Pred: positif |
|--------------------|---------------|--------------|---------------|
| Actual: negatif    | 19            | 0            | 1             |
| Actual: netral     | 4             | 0            | 1             |
| Actual: positif    | 8             | 0            | 3             |

Gambar 2. Confusion Matrix

Gambar 2 menunjukkan bahwa model memiliki kecenderungan kuat untuk memprediksi komentar ke dalam kelas negatif. Hal ini terlihat dari banyaknya data positif (8 cuitan) dan netral (4 cuitan) yang keliru diklasifikasikan sebagai negatif oleh sistem.

#### 2) Hasil Metrik Evaluasi Kuantitatif

Mengevaluasi hasil pemetaan pada matriks kebingungan, peneliti mengukur parameter akurasi, presisi, recall, serta f1-score untuk mengestimasi efektivitas sistem secara komprehensif. Perolehan akurasi final menyentuh angka 83,81%, yang menegaskan reliabilitas model dalam

mengategorikan aspirasi pengguna platform X seputar pemberhentian Patrick Kluivert. Secara mendetail, kategori negatif menyajikan hasil paling impresif dengan nilai recall 95,00% dan presisi 61,29%, membuktikan kecakapan algoritma dalam menjangkau opini bernada penolakan atau kekecewaan.

Sebaliknya, kategori positif mencatat presisi 60,00% namun dengan recall rendah (27,27%), yang menyiratkan adanya bias di mana banyak cuitan positif gagal terdeteksi dan justru teridentifikasi sebagai negatif. Kategori netral menunjukkan performa paling minim dengan skor 0,00% akibat keterbatasan sampel pembelajaran pada kelas tersebut. Adapun nilai rata-rata makro untuk presisi, recall, dan f1-score secara kolektif adalah 40,43%, 40,76%, dan 37,34%, sebagaimana terangkum dalam ilustrasi pada Gambar 3.

 Classification Report

| Class         | Precision | Recall | F1-Score | Support |
|---------------|-----------|--------|----------|---------|
| Class negatif | 61.29%    | 95.00% | 74.51%   | 20      |
| Class netral  | 0.00%     | 0.00%  | 0.00%    | 5       |
| Class positif | 60.00%    | 27.27% | 37.50%   | 11      |
| MACRO AVG     | 40.43%    | 40.76% | 37.34%   | 36      |
| WEIGHTED AVG  | 52.38%    | 61.11% | 52.85%   | 36      |

Gambar 3. Classification Report

### 3) Analisis Sampel Prediksi Kalimat

Untuk memvalidasi hasil kuantitatif, dilakukan uji coba pada beberapa sampel kalimat untuk melihat bagaimana algoritma bekerja secara tekstual. Hal ini seperti terlihat pada Tabel II berikut:

TABEL II  
CONTOH HASIL PREDIKSI KALIMAT OLEH MODEL

| No | Teks                                                                                                 | Actual  | Predicted | Status |
|----|------------------------------------------------------------------------------------------------------|---------|-----------|--------|
| 1  | Patrick Kluivert in slot out                                                                         | Positif | Negatif   | Salah  |
| 2  | Gak bosen2 aku menyuarakan kluivert out akhirnya dijawab!! Berkat smua supporter yg ikut menyuara... | Negatif | Negatif   | Benar  |
| 3  | @unmagnetism Fmd di indo bentangin KLUIVERT OUT yang besar                                           | Negatif | Negatif   | Benar  |
| 4  | Patrick Kluivert beneran out kah?                                                                    | Netral  | Negatif   | Salah  |

Kegagalan pada sampel ke-2 dan ke-3 mempertegas adanya bias model terhadap kelas mayoritas (negatif). Kata kunci seperti "out" yang sangat dominan pada data latih kelas negatif menyebabkan kalimat tanya (netral) atau harapan (positif) yang mengandung kata tersebut cenderung ditarik ke kategori negatif

oleh probabilitas likelihood. Secara keseluruhan, meskipun sistem mencapai akurasi tinggi (83,81%), terdapat ketimpangan performa antar kelas akibat ketidakseimbangan data (data imbalance). Namun, untuk tujuan pemetaan sentimen dominan, algoritma Multinomial Naïve Bayes terbukti fungsional dalam menangkap arus utama opini negatif masyarakat sepak bola Indonesia di platform X.

## IV. PENUTUP

Penelitian ini berhasil mengimplementasikan algoritma Multinomial Naive Bayes dengan pembobotan TF-IDF untuk memetakan opini publik terkait pemecatan Patrick Kluivert di platform X, menghasilkan tingkat akurasi sebesar 83,81%. Hasil analisis menunjukkan dominasi kuat sentimen negatif sebesar 59,5%, yang mencerminkan rasa kekecewaan dan penolakan mayoritas netizen terhadap keputusan manajemen Timnas Indonesia tersebut. Meskipun model bekerja sangat optimal dalam mengenali kelas mayoritas (negatif) dengan recall mencapai 95%, sistem masih mengalami kendala signifikan dalam mengklasifikasikan kelas netral akibat adanya ketidakseimbangan data (data imbalance).

Untuk pengembangan ke depan, disarankan penambahan tahap normalisasi teks guna menangani bahasa slang serta perluasan kamus stopwords khusus istilah sepak bola agar ekstraksi fitur lebih akurat. Selain itu, penggunaan teknik penyeimbangan data dan eksplorasi algoritma pembandingan seperti Support Vector Machine (SVM) atau Random Forest direkomendasikan untuk meningkatkan performa klasifikasi pada seluruh kategori sentimen secara lebih merata.

## REFERENSI

- [1] T. H. Aftoni, W. H. Sugiharto, and M. I. Ghazali, "Analisis Sentimen Terhadap Tagar #kaburajadulu di Media Sosial X Menggunakan Metode Naïve Bayes," *JUKOMTEK*, vol. 4, no. 2, 2025, doi: 10.64626/jukomtek.v4i2.443.
- [2] R. Indriati, "Analisis Sentimen Undang-Undang Cipta Kerja pada Twitter," *Malang: Penerbit Litnus*, 2024.
- [3] G. M. C. P. Man, A. A. J. Sinlae, and E. Ngaga, "Analisis Sentimen di Media Sosial X tentang IKN dengan Naïve Bayes," *JIP*, vol. 11, no. 4, pp. 417–426, 2025, doi: 10.33795/jip.v11i4.7246.
- [4] G. G. F. Djema and O. Suria, "Publik Terhadap Kebijakan Danantara Melalui Media Sosial X dengan Metode SVM dan Random Forest," *INVOTEK Polbeng - Seri Informatika*, vol. 10, no. 2, pp. 1208–1217, 2025, doi: 10.35314/rer21h75.
- [5] A. M. Jon and I. V. Papatungan, "Analisis Sentimen Pada Media Sosial Instagram Klub Persija Jakarta Menggunakan Metode Naive Bayes," *Prosiding Automata*, vol. 4, no. 1, 2023, pp. 9–16. [Online] Available At: <https://journal.uui.ac.id/AUTOMATA/article/view/26376>
- [6] M. A. Ramadhan and M. I. Wahyudin, "Analisis Sentimen Mengenai Keberhasilan Indonesia di Ajang Thomas Cup 2020 (Studi Kasus Media Sosial Twitter) Menggunakan Metode Naïve Bayes dan Decision Tree," *JTIK: Jurnal Teknologi Informasi dan Komunikasi*, vol. 6, no. 4, 2022, doi: 10.35870/jtik.v6i4.560.
- [7] N. A. Murnastiti and T. N. Fatyanosa, "Analisis Sentimen Terhadap Makanan Manis di Platform X Menggunakan TF-IDF dan Naive Bayes," *J-PTIK*, vol. 1, no. 1, pp. 1–7, 2025.
- [8] B. N. Qalbi, H. Sujaini, and R. Septiriana, "Komparasi Algoritma Naive Bayes dan Logistic Regression Terhadap Analisis Sentimen Pergantian Pelatih TIMNAS Indonesia di Media Sosial X (Twitter)," *Scientica*, vol. 4, no. 2, pp. 333–340, 2026, [Online] Available At: <https://jurnal.kolibi.id/index.php/scientica/article/view/427/413>
- [9] M. D. Fardi and M. Akbar, "Analisis Sentimen Supporter Pss Sleman Terhadap Peringkat Klub Di Twitter / X Menggunakan Naive Bayes," *Jurnal Tek. Inform.*, vol. 5, no. 3, 2025, doi: 10.58794/jekin.v5i3.1616.
- [10] N. Hadi and D. Sugiarto, "Analisis Sentimen Pembangunan IKN pada

- Media Sosial X Menggunakan Algoritma SVM, Logistic Regression dan Naïve Bayes,” *J. Inform. J. Pengemb. IT*, vol. 10, no. 1, pp. 37–49, 2025, doi: 10.30591/jpit.v10i1.7106.
- [11] Y. A. Prasetyo, E. Utami, and A. Yaqin, “Pengaruh Komposisi Split Data Terhadap Performa Akurasi Analisis Sentimen Algoritma Naïve Bayes dan SVM,” *JEECOM*, vol. 6, no. 2, pp. 382–390, 2024, doi: 10.33650/jeeecom.v4i2.
- [12] S. Adelia, *et al.*, “Analisis Sentimen Belajar Programming Pada Media Sosial Youtube Menggunakan Algoritma Klasifikasi Naïve Bayes,” *J. Inf. Technol. Ampera*, vol. 4, no. 3, pp. 254–264, 2023, [Online] Available At: <https://journal-computing.org/index.php/journal-ita/article/view/430/204>
- [13] W. B. Zulfikar, A. R. Atmadja, and S. F. Pratama, “Sentiment Analysis on Social Media Against Public Policy Using Multinomial Naïve Bayes,” *Sci. J. Informatics*, vol. 10, no. 1, pp. 25–34, 2023, doi: 10.15294/sji.v10i1.39952.
- [14] Y. Durachman, S. J. Putra, H. Nanang, and H. T. Sukmana, “Analysis Sentiment of Public Opinion on Social Media Using Naïve Bayes and TF-IDF Algorithms,” *IEEE Xplore*, 2024, doi: 10.1109/ICCIT62134.2024.10701191.
- [15] T. I. Alfawas, A. Rahim, and R. Rudiman, “Penerapan Fitur Ekstraksi TF-IDF untuk Analisis Sentimen Ulasan Game Bus Simulator Indonesia dengan Algoritma Naïve Bayes,” *Innov. J. Soc. Sci. Res.*, vol. 4, no. 5, pp. 3177–3193, 2024, doi: 10.31004/innovative.v4i5.13975
- [16] R. A. Wangsa, D. A. Kurnia, Y. A. Wijaya, and P. P. Marta, “Analisis Sentimen Bahasa Indonesia Pada Platform Metasuite Menggunakan Metode TF-IDF dan Multinomial Naïve Bayes,” vol. 02, no. 06, pp. 1–6, 2025, [Online] Available At: <https://ruangjurnal.or.id/index.php/jcsai/article/view/180>.
- [17] C. Dewi, R. C. Chen, H. J. Christanto, and F. Cauteruccio, “Multinomial Naïve Bayes Classifier for Sentiment Analysis of Internet Movie Database,” *Vietnam J. Comput. Sci.*, vol. 10, no. 4, pp. 485–498, 2023, doi: 10.1142/S2196888823500100.
- [18] F. D. Kurniawan, D. E. Ratnawati, and A. Syawli, “Analisis Sentimen Pengguna Media Sosial Twitter/X terhadap Kontribusi Megawati Hangestri Pertiwi pada Klub Bola Voli Red Sparks Korea Selatan Menggunakan Metode Multinomial Naïve Bayes dan Support Vector Machine,” *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 9, no. 8, pp. 1–10, 2025, [Online] Available At: <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/15197>.
- [19] R. Zulchamain, G. Abdurrahman, and D. Daryanto, “Analisis Sentimen Ulasan Duolingo dengan Metode Algoritma Multinomial Naïve Bayes,” *J. Inform. dan Teknol. Pendidik.*, vol. 5, no. 1, pp. 1–16, 2025, doi: 10.59395/jitp.v5i1.113.