

Analisis Sentimen Komentar Youtube Kanal Dirty Vote Menggunakan Metode Naive Bayes Classifier

1st Yuyun Khanafiyah

Teknik Informatika

Science Group

Purwokerto, Indonesia

yuyunkhanafiyah@student.telkomuniversity.ac.id

2nd Paradise

Teknik Informatika

Science Group

Purwokerto, Indonesia

paradise@telkomuniversity.ac.id

3rd Dian Kartika Sari

Teknik Informatika

Science Group

Purwokerto, Indonesia

dians@telkomuniversity.ac.id

Abstrak — Youtube merupakan media sosial yang digunakan masyarakat untuk mengekspresikan opini terhadap berbagai isu, termasuk politik. Salah satu kanal yang menjadi perhatian publik adalah *Dirty Vote*, yang memuat konten edukatif dan kritis terhadap kondisi demokrasi di Indonesia menjelang Pemilu 2024. Penelitian ini bertujuan untuk mengkaji sentimen publik terhadap konten tersebut melalui komentar pengguna. Topik ini penting karena opini publik yang terbentuk di media sosial seperti YouTube dapat mencerminkan persepsi dan respon masyarakat terhadap isu politik strategis. Namun, komentar di media sosial memiliki karakteristik teks yang tidak terstruktur, campuran bahasa, hingga sarkasme, yang menimbulkan tantangan dalam analisis otomatis. Solusi yang diterapkan adalah dengan melakukan analisis sentimen pada komentar video “*Dirty Vote – full movie*” menggunakan algoritma *Naive Bayes Classifier*. Proses mencakup *crawling* data, *preprocessing* (*cleaning*, *case folding*, *tokenisasi*, *stopword removal*, *stemming*), pelabelan menggunakan *TextBlob*, serta klasifikasi berbasis TF-IDF. Data yang dianalisis berjumlah 63 ribu komentar. Hasil pengujian menunjukkan bahwa metode *Naive Bayes Classifier* mampu mengklasifikasikan komentar dengan akurasi sebesar 71%. Untuk sentimen positif, diperoleh *precision* sebesar 70%, *recall* 91%, dan *f1-score* 79%; sementara sentimen negatif menghasilkan *precision* 76%, *recall* 42%, dan *f1-score* 54%. Penelitian ini memberikan gambaran umum mengenai respon publik terhadap konten politik di media sosial.

Kata kunci— Analisis Sentimen, Youtube, *Naive Bayes Classifier*, *Dirty Vote*, Komentar, Politik

I. PENDAHULUAN

Media sosial seperti YouTube menjadi ruang publik modern yang memfasilitasi penyebaran opini dan diskusi terhadap berbagai isu, termasuk isu politik[1]. Di Indonesia, kanal YouTube *Dirty Vote* muncul sebagai salah satu platform yang menyuarakan kritik terhadap kondisi demokrasi menjelang Pemilu 2024. Video pertamanya, berjudul “*Dirty Vote – full movie*”, menyajikan pandangan tiga pakar hukum tata negara terhadap penyalahgunaan kekuasaan yang dianggap merusak demokrasi. Respons publik terhadap video ini sangat tinggi, tercermin dari

puluhan ribu komentar yang menunjukkan beragam pandangan dan emosi.

Fenomena ini mencerminkan pentingnya analisis terhadap opini publik di ruang digital, mengingat media sosial telah menjadi sumber data yang merepresentasikan persepsi masyarakat secara luas. Namun, komentar di media sosial memiliki karakteristik tidak terstruktur, menggunakan bahasa campuran, dan sering kali mengandung sarkasme atau ekspresi informal, sehingga menyulitkan analisis secara manual. Oleh karena itu, pendekatan otomatis berbasis machine learning diperlukan untuk mengklasifikasikan sentimen publik secara efektif.

Salah satu teknik yang sering dipakai dalam klasifikasi teks adalah algoritma *Naive Bayes Classifier* (NBC). NBC dikenal karena kesederhanaannya, efisiensi komputasi, dan kinerjanya yang cukup kompetitif dalam tugas klasifikasi berbasis teks[2]. Sejumlah penelitian terdahulu menunjukkan bahwa NBC mampu menghasilkan tingkat akurasi yang tinggi dalam analisis sentimen, baik di bidang politik, kebijakan publik, hingga konten hiburan.

Penelitian ini bertujuan untuk menganalisis sentimen komentar terhadap video “*Dirty Vote – full movie*” dengan memanfaatkan algoritma *Naive Bayes Classifier*. Tahap analisis meliputi tahap *crawling* data komentar dari YouTube, pra-pemrosesan teks seperti pembersihan data, tokenisasi, *stopword removal*, dan *stemming*, hingga pelabelan menggunakan *TextBlob* dan klasifikasi berbasis TF-IDF.

Dengan mengkaji distribusi komentar positif dan negatif serta mengevaluasi kinerja model klasifikasi, penelitian ini diharapkan memberikan gambaran umum mengenai persepsi masyarakat terhadap konten politik yang bersifat kritis. Hasil penelitian ini juga diharapkan mampu memberikan kontribusi dalam pengembangan teknik analisis sentimen serta memperkaya literatur terkait penggunaan machine learning dalam pemetaan opini publik di media sosial.

II. KAJIAN TEORI

A. Youtube

YouTube merupakan platform berbasis internet milik Google yang memiliki berbagai fitur seperti mengunggah dan menampilkan video maupun animasi, serta pengguna dapat membagikan konten mereka. Selain itu Youtube dapat diakses oleh pengguna dari seluruh dunia. Kelebihan Youtube antara lain dapat mengunggah video, mengunduh video, melakukan *streaming* dan memilih kualitas video sesuai keinginan saat menonton. Kelemahan dari Youtube biasanya disalah gunakan oleh pengguna seperti untuk ujaran kebencian, pornografi, spam, serta banyak informasi yang tidak sesuai kebenarannya[3].

B. API v3

YouTube Data API v3 merupakan layanan publik yang memungkinkan pengembang mengakses data YouTube seperti video, channel, dan komentar secara terprogram[4]. API ini memungkinkan sistem berinteraksi tanpa harus memahami detail internal satu sama lain. Untuk menggunakannya, pengguna perlu mengaktifkan layanan melalui Google Cloud Console dan memperoleh API key sebagai autentikasi akses data[5].

C. Analisis Sentimen

Analisis sentimen atau *opinion mining* merupakan metode pengolahan bahasa alami (*Natural Language Processing/NLP*) untuk mengidentifikasi opini atau emosi yang terkandung dalam teks[6]. Tujuan utamanya adalah mengelompokkan teks ke dalam kategori sentimen, seperti positif, negatif, atau netral. Dalam penelitian ini, klasifikasi difokuskan pada dua label utama yaitu positif dan negatif. Analisis sentimen banyak diterapkan dalam berbagai domain seperti pemasaran, politik, dan media sosial [7].

D. Pra-pemrosesan Teks

Teks komentar di media sosial sering kali tidak terstruktur, mengandung bahasa informal, emoji, hingga campuran bahasa. Oleh karena itu, perlu dilakukan tahapan pra-pemrosesan agar data dapat diolah secara efektif. Langkah-langkah utama dalam preprocessing adalah:

- *Cleaning*: Menghapus URL, simbol, emoji, dan teks tidak relevan lainnya.
- *Case Folding*: Mengubah seluruh huruf menjadi huruf kecil.
- *Normalisasi*: Mengubah istilah yang tidak formal menjadi bentuk baku (misal: “ga” menjadi “tidak”).
- *Tokenisasi*: Membedakan kalimat menjadi kata-kata (*token*).
- *Stopword Removal*: Menghapus kata-kata umum yang tidak memiliki makna penting (seperti “yang”, “dan”).
- *Stemming*: Mengembalikan kata ke bentuk dasarnya (misal: “berlari” menjadi “lari”) [2].

E. Pelabelan Data

Proses pelabelan data merupakan tahap pemberian kategori atau label pada data. Proses ini bisa dilakukan secara manual atau otomatis dengan memanfaatkan algoritma pembelajaran mesin[8]. Pelabelan dilakukan untuk menentukan kelas sentimen dari masing-masing komentar. Dalam penelitian ini digunakan pustaka *TextBlob*, yang secara otomatis menghitung nilai polaritas kalimat. Jika polaritas > 0 maka dikategorikan sebagai positif, dan jika < 0

dikategorikan sebagai negatif. Komentar dengan polaritas nol diberi label secara acak untuk menjaga distribusi data.

F. TF-IDF

TF-IDF (*Term Frequency Inverse Document Frequency*) merupakan teknik pembobotan kata yang berfungsi untuk menilai seberapa signifikan suatu kata dalam sebuah dokumen[9]. Nilai TF menunjukkan frekuensi kemunculan kata itu muncul dalam dokumen, sedangkan IDF menunjukkan sejauh mana kata tersebut jarang muncul di keseluruhan dokumen. Kata-kata yang memiliki nilai TF-IDF tinggi dianggap lebih relevan dalam membedakan dokumen antar kategori[10].

G. Algoritma Naive Bayes Classifier

Naive Bayes Classifier merupakan algoritma klasifikasi yang bersifat probabilistik dan memanfaatkan Teorema Bayes dengan anggapan bahwa setiap fitur bersifat independen satu sama lain[11]. Algoritma ini menghitung peluang sebuah data tergolong dalam kategori tertentu dengan melihat seberapa sering fitur muncul dalam data pelatihan. Meskipun sederhana dan mengandalkan asumsi kuat, metode ini terbukti berhasil dalam pekerjaan mengklasifikasikan teks seperti dalam analisis sentimen[5]. Dasar matematisnya dinyatakan sebagai berikut:

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)} \quad (1)$$

Keterangan :

$P(A|B)$: peluang hipotesis A setelah mengamati data B (posterior)

$P(B|A)$: peluang data B jika hipotesis A terbukti benar (likelihood)

$P(A)$: peluang awal untuk hipotesis A (*prior*)

$P(B)$: peluang untuk data B (*evidence*).

Model ini dipilih dalam penelitian karena kecepatan, kemudahan implementasi, serta akurasi yang baik meskipun pada data teks yang tidak terstruktur seperti komentar media sosial.

H. Confusion Matrix

Confusion matrix dipakai untuk menilai efektivitas model klasifikasi dengan cara membandingkan hasil yang prediksi dengan label yang sebenarnya. Evaluasi ini mencakup empat komponen: *true positive* (TP), *true negative* (TN), *false positive* (FP), dan *false negative* (FN). Dari elemen-elemen ini, dapat dihitung berbagai metrik seperti akurasi, *precision*, dan *recall*. Akurasi menunjukkan persentase prediksi yang tepat, *precision* mengukur ketepatan prediksi positif, sedangkan *recall* menunjukkan seberapa baik model mengenali data positif. Ketiga metrik ini memberikan gambaran menyeluruh tentang efektivitas model dalam mengklasifikasikan sentimen [12].

I. Python dan Google Colab

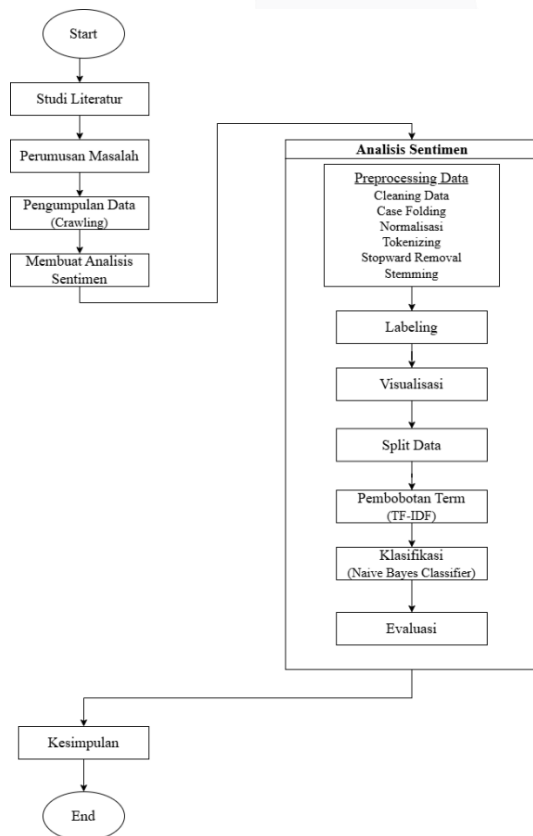
Python merupakan bahasa pemrograman yang memiliki tingkat tinggi yang diciptakan oleh Guido van Rossum pada tahun 1991. Dikenal karena sintaksisnya yang sederhana dan mudah dipahami, *Python* memungkinkan berbagai cara dalam berpikir mengenai pemrograman, termasuk cara berorientasi objek dan fungsional. *Python* banyak digunakan dalam pengembangan web, automasi, analisis data, kecerdasan buatan, dan pengolahan bahasa alami. Dukungan pustaka yang luas seperti *Pandas*, *NLTK*, *Scikit-learn*, dan

TensorFlow menjadikan *Python* pilihan utama dalam penelitian dan pengembangan teknologi modern[13].

Untuk mendukung pemrograman *Python* secara praktis, Google menyediakan Google Colaboratory (Colab), yaitu platform *cloud-based* yang memungkinkan pengguna menjalankan kode *Python* langsung dari browser tanpa instalasi tambahan. Google Colab mendukung penggunaan GPU/TPU, menyimpan file langsung ke Google Drive, dan mendukung kolaborasi secara real-time. Platform ini sangat cocok untuk eksperimen machine learning dan analisis data berskala besar seperti yang digunakan dalam penelitian ini[6].

III. METODE

Penelitian ini merupakan penelitian eksperimen yang bertujuan untuk menilai seberapa efektif algoritma *Naïve Bayes Classifier* dalam mengklasifikasikan sentimen komentar pengguna mengenai video “*Dirty Vote – full movie*” yang ada di platform YouTube *Dirty Vote*. Penelitian dilakukan melalui serangkaian tahapan terstruktur yang mencakup studi literatur, perumusan masalah, pengumpulan data, pra-pemrosesan teks, pelabelan sentimen, ekstraksi fitur, pelatihan model klasifikasi, serta evaluasi performa model. Hubungan yang dikaji dalam penelitian ini mencakup pengaruh data masukan berupa komentar pengguna YouTube (variabel X) terhadap hasil klasifikasi sentimen (variabel Y), yang terbagi dalam kategori positif dan negatif. Berikut gambar diagram alir tahapan penelitian ini :



GAMBAR 1
DIAGRAM ALIR PENELITIAN

Langkah pertama yang dilakukan dalam penelitian ini adalah studi literatur, yang bertujuan untuk memperoleh

landasan teoritis dan memperkuat pemahaman terkait topik yang dikaji. Sumber referensi yang digunakan diperoleh dari jurnal dan artikel ilmiah yang relevan, diakses melalui platform terpercaya seperti Google Scholar, ScienceDirect, dan RabbitResearch. Batasan waktu publikasi ditetapkan dalam rentang lima tahun terakhir untuk memastikan aktualitas dan relevansi referensi.

Setelah studi literatur, peneliti merumuskan permasalahan dengan mengidentifikasi isu yang berkembang dari video “*Dirty Vote – full movie*” melalui komentar-komentar yang ditinggalkan oleh pengguna pada kanal YouTube *Dirty Vote*. Perumusan masalah juga diperkuat dengan informasi dari berbagai media daring yang kredibel, seperti CNN Indonesia, CNBC Indonesia, dan KumparanTech, yang banyak membahas respons publik terhadap film dokumenter tersebut. Permasalahan utama yang akan diteliti dalam penelitian ini adalah cara mengklasifikasikan sentimen (positif dan negatif) dari komentar-komentar tersebut, serta mengevaluasi sejauh mana algoritma *Naïve Bayes Classifier* mampu mengotomatisasi proses klasifikasi sentimen secara akurat.

Proses pengumpulan data dilakukan dengan memanfaatkan YouTube Data API v3 untuk mengambil komentar dari video “*Dirty Vote – full movie*” yang diunggah pada kanal resmi *Dirty Vote*. Pengambilan data dilakukan pada tanggal 20 Maret 2024, dengan total data komentar yang berhasil dikumpulkan mencapai 63.000 entri. Data disimpan dalam format CSV dan digunakan sebagai bahan utama dalam proses analisis sentimen.

Setelah data terkumpul, dilakukan pra-pemrosesan untuk membersihkan dan menyiapkan data teks sebelum dianalisis. Langkah-langkah pra-pemrosesan meliputi *cleaning*, *case folding*, normalisasi, *tokenizing*, *stopword removal*, dan *stemming*. Proses ini dilakukan dengan bantuan pustaka *Python* seperti NLTK dan Sastrawi. Data yang telah dibersihkan, komentar diterjemahkan ke dalam bahasa Inggris menggunakan pustaka Google Translate agar dapat diproses oleh pustaka *TextBlob* untuk pelabelan sentimen. Hasil pelabelan ini kemudian divisualisasikan menggunakan pustaka *WordCloud*, *Matplotlib*, dan *NumPy* untuk menampilkan distribusi kata dan proporsi sentimen secara visual dan informatif.

Sebelum proses ekstraksi data dilakukan, data dipisahkan menjadi dua bagian, yaitu 90% untuk pelatihan dan 10% untuk pengujian. Pemisahan ini dilakukan memanfaatkan fungsi *train_test_split* yang ada dalam pustaka *Scikit-learn*. Rasio tersebut dipilih untuk memastikan cakupan pembelajaran yang luas dan evaluasi yang akurat terhadap model yang dikembangkan. Ekstraksi fitur dilakukan dengan pendekatan *Term Frequency-Inverse Document Frequency* (TF-IDF). Metode ini dipakai untuk mengoversi data teks menjadi bentuk angka yang memuat informasi penting dari setiap kata dalam dokumen, sehingga meningkatkan akurasi proses klasifikasi.

Setelah data teks direpresentasikan dalam bentuk vektor numerik menggunakan metode TF-IDF, langkah selanjutnya adalah penerapan algoritma *Naïve Bayes Classifier* untuk proses klasifikasi. Algoritma ini dipilih karena terkenal dengan keunggulannya dalam hal kemudahan penerapan, kecepatan proses, dan kemampuan yang baik dalam mengelola data teks dalam jumlah besar seperti komentar YouTube. Model dilatih dengan menggunakan data pelatihan dan kemudian diuji dengan data pengujian untuk menilai

kinerjanya. Evaluasi dilakukan dengan menghitung akurasi menggunakan *confusion matrix* dan sejumlah metrik lainnya seperti *precision*, *recall*, dan *f1-score*. Hasil evaluasi ini memberikan gambaran menyeluruh mengenai kemampuan model dalam mengelompokkan komentar ke dalam kategori sentimen positif dan negatif secara tepat.

Tahap terakhir dari penelitian ini adalah menngumpulkan kesimpulan berdasarkan hasil klasifikasi dan evaluasi model. Kesimpulan tersebut diharapkan dapat memberikan pemahaman mengenai persepsi publik terhadap konten politik di media sosial serta berfungsi sebagai acuan untuk studi-studi mendatang di area analisis sentimen dan pengolahan data opini publik.

IV. HASIL DAN PEMBAHASAN

Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap video “*Dirty Vote – full movie*” dengan mengklasifikasikan komentar pengguna YouTube ke dalam dua kategori: positif dan negatif. Proses penelitian meliputi langkah-langkah pengumpulan data, pembersihan data, pelabelan data, pembobotan fitur, klasifikasi menggunakan *Naïve Bayes Classifier*, serta evaluasi performa model.

A. Pengumpulan dan Pra-Pemrosesan Data

Data dikumpulkan menggunakan YouTube Data API v3 dari kanal *Dirty Vote* dengan total 63.000 komentar yang dikumpulkan pada tanggal 20 Maret 2024. Setiap entri terdiri dari atribut waktu unggah, nama pengguna, isi komentar, dan jumlah *like*. Sampel data ditampilkan pada Tabel 1(a).

TABEL 1
(A)

<i>publishAt</i>	<i>authorDisplay Name</i>	<i>textDisplay</i>	<i>Like Count</i>
2024-02-26T23:23:48Z	@synechisebrosta5200	Mantap...	1
2024-02-27T16:45:58Z	@zenenff7273	Di tunggu part II nya pasca pemilu ,, rekap suara, deklarasi awal, pengangkatan menteri, hak angket, waaaah...banyak deh.. ayo bang observasi dulu.. biar JD dokumen untuk anak2 cucu kita nanti	3
2024-02-27T17:25:14Z	@khairudinudin6353	ðŸ˜ˆ,ðŸ˜ˆ,ðŸ˜ˆ,cemen kritikan gue dihapus wkwkwk cemennnnn	8
2024-02-27T19:53:10Z	@resepemase	ini film dg actor dan scenario terburuk ,,,, niat mau nambah suara malah anjlok di bawah 20 % saja ,,,	13
2024-03-20T02:03:52Z	@fredisuprianto3621	Nanti kalo 2 putaran 80% , kejang"	0

Selanjutnya, komentar dikonversi menjadi teks bersih dengan tahapan *cleaning*, *case folding*, *normalisasi*, *tokenizing*, *stopword removal*, dan *stemming*. Tabel berikut menunjukkan hasil transformasi teks pada tiap tahap.

TABEL 2
(A)

Tahap	<i>textDisplay</i>	Hasil
<i>Cleaning</i>	@ Ini film dg actor dan scenario terburuk ,,,, niat mau nambah	Ini film dg actor dan scenario terburuk brniat mau nambah suara malah Anjlok di bawah saja
<i>Case Folding</i>	Ini film dg actor dan scenario terburuk brniat mau nambah suara malah Anjlok di bawah saja	ini film dg actor dan scenario terburuk brniat mau nambah suara malah anjlok di bawah saja
Normalisasi	ini film dg actor dan scenario terburuk brniat mau nambah suara malah anjlok di bawah saja	ini film dengan aktor dan skenario terburuk berniat mau menambah suara malah anjlok di bawah saja
<i>Tokenizing</i>	ini film dengan aktor dan skenario terburuk berniat mau menambah suara malah anjlok di bawah saja	['ini', 'film', 'dengan', 'aktor', 'dan', 'skenario', 'terburuk', 'brniat', 'mau', 'menambah', 'suara', 'malah', 'anjlok', 'di', 'bawah', 'saja']
<i>Stopword Removal</i>	['ini', 'film', 'dengan', 'aktor', 'dan', 'skenario', 'terburuk', 'brniat', 'mau', 'menambah', 'suara', 'malah', 'anjlok', 'di', 'bawah', 'saja']	film aktor skenario terburuk berniat mau menambah suara malah anjlok bawah
<i>Stemming</i>	film aktor skenario terburuk berniat mau menambah suara malah anjlok bawah	film aktor skenario buruk niat mau tambah suara malah anjlok bawah

B. Pelabelan Sentimen dan Visualisasi

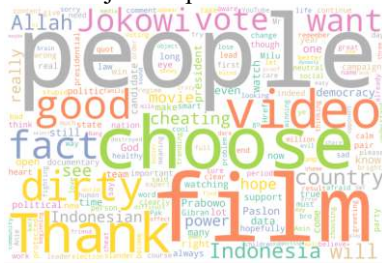
Komentar yang telah diproses kemudian diterjemahkan ke dalam bahasa Inggris melalui Google Translate API, lalu dilakukan pelabelan sentimen menggunakan pustaka *TextBlob*. Komentar dengan nilai polaritas lebih dari nol diklasifikasikan sebagai sentimen positif, sedangkan komentar dengan nilai polaritas kurang dari nol diklasifikasikan sebagai sentimen negatif. Untuk komentar yang memiliki nilai polaritas tepat nol, label diberikan secara acak ke dalam salah satu dari dua kategori, yaitu positif atau negatif. Pendekatan ini digunakan untuk menjaga keseimbangan distribusi data, mengingat penelitian ini hanya memfokuskan pada dua kelas sentimen.

Hasil pelabelan menunjukkan bahwa dari total 58.091 komentar yang berhasil diproses, sebanyak 34.907 komentar (60,1%) dikategorikan sebagai sentimen positif, sedangkan 23.184 komentar (39,9%) dikategorikan sebagai sentimen negatif. Persentase yang dominan pada kategori positif mencerminkan bahwa mayoritas penonton memberikan respons yang mendukung atau setidaknya sejalan dengan isi video yang ditayangkan. Berikut Tabel *labeling* data.

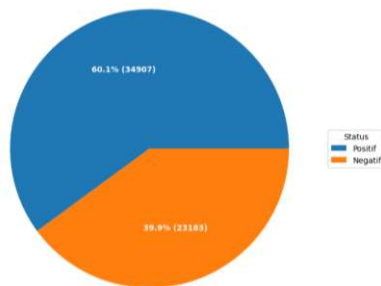
<i>text preprocessing</i>	<i>Sentiment</i>
mantap	Positif
tunggu part ii nya pasca pemilu rekap suara deklarasi awal angkat menteri hak angket banyak ayo bang observasi dulu biar jadi dokumen anak cucu nanti	Positif
cemen kritik saya hapus cemen	Negatif
film aktor skenario buruk berniat mau tambah suara malah anjlok bawah	Negatif
kalo putar kejang	Negatif

Setelah pelabelan, data divisualisasikan untuk memperjelas pola dan distribusi komentar. Salah satu visualisasi yang digunakan adalah *WordCloud*, yang menampilkan kata-kata berdasarkan frekuensi

kemunculannya. Kata yang sering muncul ditampilkan lebih besar, sehingga membantu mengidentifikasi tema dominan dalam komentar terhadap video "*Dirty Vote – full movie*". Hasil visualisasi ditunjukkan pada Gambar 2.



Visualisasi selanjutnya menggunakan diagram lingkaran (*Pie Chart*) untuk menampilkan proporsi sentimen positif dan negatif. Hasilnya menunjukkan dominasi komentar positif, sejalan dengan hasil pelabelan. Distribusi sentimen ini ditampilkan pada Gambar 3.



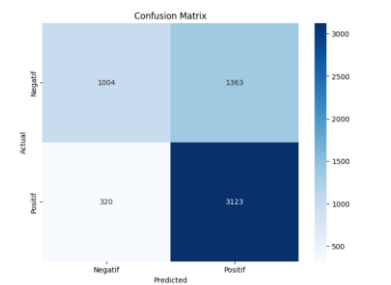
C. Pembagian Data dan Pembobotan Fitur

Tahap selanjutnya adalah membagi data menjadi 90% data pelatihan (52.281 data) dan 10% data pengujian (5.810 data). Rasio ini dipilih berdasarkan uji coba beberapa skenario, seperti 50:50 hingga 80:20. Hasilnya, rasio 90:10 memberikan akurasi tertinggi. Dengan data pelatihan yang lebih banyak, model *Naïve Bayes* mampu mempelajari pola sentimen secara lebih optimal, tanpa mengorbankan kualitas evaluasi pada data uji.

Setelah itu, dilakukan pembobotan fitur menggunakan TF-IDF untuk mengubah komentar teks diubah menjadi bentuk vektor numerik. Teknik ini menentukan bobot dengan mempertimbangkan seberapa sering kata muncul dalam dokumen (TF) dan seberapa tidak biasanya kata tersebut muncul dalam seluruh dataset (IDF). Kata yang jarang dan spesifik akan memiliki bobot tinggi, sementara kata umum akan memiliki bobot rendah. Proses ini dilakukan menggunakan *TfidfVectorizer* dari *Scikit-learn*.

D. Klasifikasi dan Evaluasi

Tahap klasifikasi dilakukan menggunakan algoritma *Multinomial Naïve Bayes* (MNB) yang efektif untuk data teks diskrit seperti komentar YouTube. Meskipun MNB mengasumsikan independensi antar fitur, algoritma ini terbukti handal dan efisien dalam tugas klasifikasi teks. Model dilatih menggunakan data yang telah melalui proses pembobotan TF-IDF, kemudian diuji untuk memprediksi sentimen komentar menjadi dua kategori, yakni positif dan negatif. Penilaian performa model dilakukan dengan memanfaatkan *confusion matrix* dan empat metrik penting yaitu akurasi, presisi, *recall*, serta *f1-score*. Berikut hasil menggunakan *confusion matrix*.



	precision	recall	f1-score	support
0	0.76	0.42	0.54	2367
1	0.70	0.91	0.79	3443
accuracy			0.71	5810
macro avg	0.73	0.67	0.67	5810
weighted avg	0.72	0.71	0.69	5810

Berdasarkan hasil tersebut, model terbukti lebih akurat dalam mendeteksi komentar positif, yang ditunjukkan oleh nilai *recall* yang tinggi. Sebaliknya, kemampuan model dalam mengenali komentar negatif masih rendah, dengan *recall* hanya sebesar 42%. Ketidakseimbangan ini dapat disebabkan oleh distribusi data yang lebih condong ke sentimen positif serta keterbatasan asumsi independensi fitur pada algoritma *Naïve Bayes*. Meski demikian, nilai *precision* yang cukup tinggi pada kedua kelas menunjukkan bahwa prediksi yang dihasilkan masih cukup dapat diandalkan. Dengan akurasi keseluruhan sebesar 71%, model ini cukup efektif sebagai alat bantu dalam memetakan opini publik terhadap konten politik di media sosial.

V. KESIMPULAN

Hasil penelitian menunjukkan bahwa algoritma *Naive Bayes Classifier* dapat dimanfaatkan dengan baik untuk menganalisis sentimen komentar di YouTube pada video “*Dirty Vote – Full Movie*”. Dari 58.091 komentar yang dianalisis, sebanyak 60,1% tergolong positif, menunjukkan respons publik yang cenderung mendukung konten tersebut. Model menghasilkan akurasi sebesar 71%, dengan kinerja lebih baik pada klasifikasi sentimen positif (*recall* 91%) dibandingkan negatif (*recall* 42%). Hal ini menunjukkan bahwa meskipun model cukup baik untuk mendeteksi komentar positif, masih diperlukan pengembangan lebih lanjut untuk meningkatkan akurasi klasifikasi komentar negatif. Oleh karena itu, peneliti selanjutnya dapat mempertimbangkan perbandingan dengan metode *machine learning* lain, seperti *Support Vector Machine* (SVM), *Random Forest*, atau algoritma *deep learning* seperti LSTM, yang lebih baik dalam menangani ketidakseimbangan data dan hubungan antar fitur yang lebih kompleks. Disamping itu, penyempurnaan model *Naive Bayes* dapat dilakukan dengan teknik penyeimbangan data atau penerapan metode pengolahan teks lanjutan seperti TF-IDF atau *word embeddings* juga dapat meningkatkan akurasi dan keseimbangan dalam mengklasifikasikan komentar YouTube secara lebih tepat.

REFERENSI

- [1] Christie Passaris, “15 Ide Channel Youtube Untuk Menginspirasi Anda,” *Clipchamp*, 2022. Accessed: May

- 17, 2024. [Online]. Available: <https://clipchamp.com/id/blog/ide-channel-youtube/>
- [2] S. Mulyani And R. Novita, "Implementation Of The Naive Bayes Classifier Algorithm For Classification Of Community Sentiment About Depression On Youtube," *Jurnal Teknik Informatika (Jutif)*, Vol. 3, No. 5, Pp. 1355–1361, Oct. 2022, Doi: 10.20884/1.Jutif.2022.3.5.374.
- [3] I. P. Hendika Permana, "Analisis Rasio Pada Akun Youtube Untuk Penelitian Kualitatif Menggunakan Metode Eksploratif," *Jurnal Ilmiah Media Sisfo*, Vol. 15, No. 1, Pp. 40–48, Apr. 2021, Doi: 10.33998/Mediasisfo.2021.15.1.970.
- [4] D. E. Saputra And A. R. Isnain, "Implementasi Algoritma Convolutional Neural Network Untuk Analisis Sentimen Bacapres 2024 Pada Kolom Komentar Youtube Mata Najwa," *Jipi (Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika)*, Vol. 9, No. 3, Pp. 1431–1441, Aug. 2024, Doi: 10.29100/Jipi.V9i3.5420.
- [5] A. H. Nurridha, W. Hidayat, And A. Erfina, "Sistemasi: Jurnal Sistem Informasi Analisis Sentimen Isu Dihapusnya Kurikulum Merdeka Menggunakan Algoritma Naïve Bayes Classifier." [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [6] S. T. , M. Dr. Dra. D. K. M. K. L. F. M. H. S. Kom. , M. K. S. A. R. S. Yessy Asri, *Machine Learning & Deep Learning: Analisis Sentimen Menggunakan Ulasan Pengguna Aplikasi*. Ponorogo: Uwais Inspirasi Indonesia, 2024.
- [7] D. C. Ardhi And D. P. Sari, "Sentiment Analysis Of Youtube Comments: Potential Indonesian Presidential Election Candidates," *International Journal Of Computer Applications Technology And Research*, Pp. 451–456, Nov. 2022, Doi: 10.7753/Ijcatr1112.1010.
- [8] Mit. , Ph. D. , R. G. A. M. Kom. , Ph. D. , G. S. Kom. , M. Kom. , D. A. St. , M. Eng. Prihandoko, *Memahami Konsep Dan Implementasi Machine Learning*. Jambi: Pt. Sonpedia Publishing Indonesia, 2024.
- [9] D. Farah Zhafira, B. Rahayudi, And P. Korespondensi, "Analisis Sentimen Kebijakan Kampus Merdeka Menggunakan Naive Bayes Dan Pembobotan Tf-Idf Berdasarkan Komentar Pada Youtube," 2021.
- [10] R. Wati, S. Ernawati, And H. Rachmi, "Pembobotan Tf-Idf Menggunakan Naïve Bayes Pada Sentimen Masyarakat Mengenai Isu Kenaikan Bipih," *Jurnal Manajemen Informatika (Jamika)*, Vol. 13, No. 1, Pp. 84–93, Apr. 2023, Doi: 10.34010/Jamika.V13i1.9424.
- [11] E. T. K. Kusriani, *Algoritma Data Mining*. Yogyakarta: C.V Andi Offset, 2009.
- [12] A. Wahid And G. Saputri, "Analisis Sentimen Komentar Youtube Tentang Relawan Patwal Ambulance Menggunakan Algoritma Naïve Bayes Dan Decision Tree," *Jurnal Sistem Komputer Dan Informatika (Json)*, Vol. 4, No. 2, P. 319, Dec. 2022, Doi: 10.30865/Json.V4i2.4941.
- [13] Adawiyah Ritonga And Yahfizham Yahfizham, "Studi Literatur Perbandingan Bahasa Pemrograman C++ Dan Bahasa Pemrograman Python Pada Algoritma Pemrograman," *Jurnal Teknik Informatika Dan Teknologi Informasi*, Vol. 3, No. 3, Pp. 56–63, Nov. 2023, Doi: 10.55606/Jutiti.V3i3.2863.