



Analisis Sentimen Terhadap Boikot Produk Israel Menggunakan Algoritma Naive Bayes Dan SMOTE

Putra Ganda Dewata¹, Azzar Rizky², Hafiz Irsyad³

^{1,2,3}Program Studi Informatika, Fakultas Ilmu Komputer dan Rekayasa, Universitas Multi Data Palembang, Kota Palembang, Indonesia

¹putraganda.dewata@mhs.mdp.ac.id, ²azzar.arza@mhs.mdp.ac.id, ³hafizirsyad@mdp.ac.id

Abstrak

The conflict between Palestine and Israel, which has been ongoing for approximately six decades, has triggered a wide range of reactions in various countries. The response from the Indonesian public has included a boycott of products affiliated with Israel as a form of solidarity and moral protest against this prolonged conflict. This study aims to analyze the sentiment of the Indonesian public towards this boycott action on the YouTube Shorts video-sharing platform. Sentiment analysis in this study uses the Naive Bayes algorithm and text preprocessing techniques, specifically the SMOTE technique, with a dataset of 2433 public comments. The final results indicate that the Naive Bayes algorithm has a high level of effectiveness in predicting sentiment in text comments, with an accuracy of 94%. This level of accuracy was achieved by training the algorithm model with a data split of 60% training data and 40% testing data. Additionally, the study found that the sentiment of the Indonesian public towards the boycott of Israeli products tends to be negative.

Kata Kunci: Boycott Action, Indonesian Public, Naive Bayes, Palestine-Israel Conflict, Sentiment, SMOTE.

1. Pendahuluan

Konflik yang terjadi antara Palestina dan Israel sudah terjadi sejak lama, kurang lebih terjadi sejak kurang lebih 6 dekade lalu, jika dilihat dari pendirian resmi negara Israel pada 15 Mei 1948. Awal jatuhnya negara Palestina bermula pada abad ke-19 Masehi, yang disebabkan karena kelemahan pemerintahan yang dipimpin oleh Kerajaan Turki Uthmaniyyah pada masa itu [1].

Konflik ini hingga sekarang tidak memiliki titik cerah untuk kedua pihak. Palestina bersikeras untuk mempertahankan wilayahnya, dan tidak akan merelakan wilayahnya dikuasai oleh Israel. Israel tetap beranggapan bahwa tanah Palestina yang saat ini diduduki oleh bangsa Arab (Palestina) adalah tanah milik leluhur mereka [2], dan didukung oleh Amerika Serikat yang terus mendukung Israel dengan menggunakan hak vetonya dalam sidang keamanan PBB (Perserikatan Bangsa Bangsa) untuk menolak resolusi perdamaian terkait konflik Israel dan Palestina [3].

Dan yang terbaru, konflik Israel dan Palestina kembali terjadi yaitu pada tanggal 7 Oktober 2023, dimana terjadi eskalasi besar dalam konflik antara Israel dan Palestina. Pada hari tersebut, kelompok Hamas, yang menguasai Jalur Gaza, meluncurkan serangan roket besar-besaran ke wilayah Israel. Serangan ini dibalas dengan serangan udara oleh militer Israel ke berbagai target di Gaza. Eskalasi ini menyebabkan korban jiwa

dan kerusakan infrastruktur di kedua belah pihak, terutama paling signifikan pada wilayah Palestina.

Konflik tersebut memicu reaksi beragam yang luas di berbagai negara, tidak terkecuali yaitu Indonesia. Salah satu reaksi yang dilakukan oleh masyarakat Indonesia yaitu Aksi Boikot Produk yang terafiliasi dengan negara Israel. Aksi boikot tersebut merupakan bentuk solidaritas dan protes moral terhadap konflik berkepanjangan antara Palestina dan Israel. Aksi tersebut menimbulkan dinamika opini diantara masyarakat Indonesia dan mencakup beragam sudut pandang, baik yang mendukung maupun menentang aksi tersebut.

Salah satu media yang digunakan oleh masyarakat Indonesia dalam menyampaikan opini dan reaksi dari konflik tersebut adalah platform *online video sharing* YouTube. YouTube merupakan platform berbagi video yang dibuat oleh Google. Hingga per bulan Oktober tahun 2023, pengguna YouTube di Indonesia mencapai 139 juta pengguna, dan menjadi negara keempat dengan pengguna YouTube terbanyak di dunia.

Tingginya jumlah pengguna tersebut dapat dimanfaatkan untuk dilakukan analisis mengenai pendapat pengguna pada aksi tersebut. Data yang diambil yaitu data komentar pengguna pada video di YouTube yang relevan dengan aksi boikot tersebut. Data kemudian diolah dan dianalisis untuk dilakukan analisis sentimen. Analisis sentimen bertujuan untuk

memahami apakah opini atau sentimen yang diungkapkan dalam teks cenderung positif atau negatif [4].

Penelitian mengenai analisis sentimen terhadap aksi boikot produk yang terafiliasi dengan negara Israel telah dilakukan oleh Susilawati, A. T., Lestari, N. A., dan Nina, P. A. pada tahun 2024. Penelitian ini membahas analisis sentimen masyarakat Indonesia terkait konflik Israel-Palestina melalui platform media sosial Twitter, dengan fokus pada boikot produk Israel. Penelitian ini menggunakan *tools* Orange dan *klasifikasi Naive Bayes* untuk menganalisis lebih dari 300 dataset tweet yang diambil melalui proses *scrapping*. Hasilnya menunjukkan bahwa mayoritas masyarakat cenderung mendukung boikot produk Israel, dengan tingkat akurasi klasifikasi *Naive Bayes* mencapai 95% [5].

Pada penelitian yang sama mengenai analisis sentimen terhadap aksi boikot produk yang terafiliasi dengan negara Israel telah dilakukan oleh Az-haari, N. F., Juardi, D., dan Jamaludin, A. pada tahun 2024. Penelitian ini membahas analisis sentimen publik terhadap boikot produk yang mendukung Israel, dengan fokus pada produk Amerika seperti Starbucks yang terkena dampak setelah konflik antara Israel dan Hamas pecah pada 7 Oktober 2023. Hasil yang didapat adalah akurasi tertinggi terdapat pada model dengan pembagian data 90% training dan 10% testing dengan akurasi sebesar 71%. [6].

Pada penelitian lain juga dilakukan oleh Hayurian, L. A., & Hendrastuty, N di tahun 2024 yaitu membahas mengenai analisis sentimen publik terhadap aksi boikot produk Israel di platform sosial media Twitter menggunakan *Naive Bayes Algorithm* dan *Support Vector Machine (SVM)*. Hasil dari penelitian tersebut menunjukkan akurasi *Naive Bayes Algorithm* sebesar 85% dan *Support Vector Machine (SVM)* sebesar 84% dengan pembagian 80% *data training* dan 20% *data testing* [7].

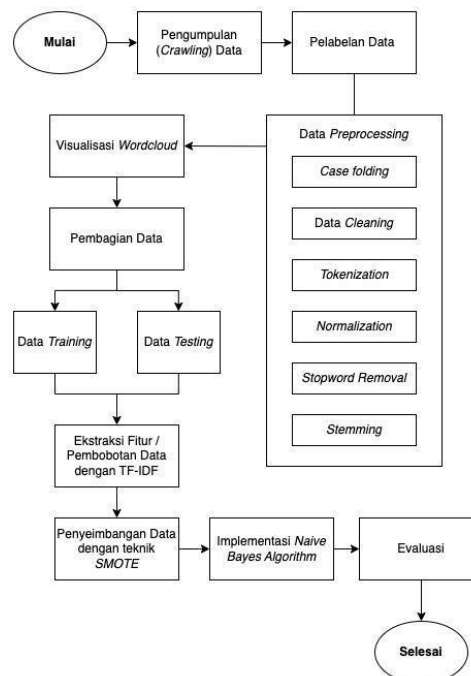
Berdasarkan dari penelitian-penelitian tersebut, algoritma yang akan digunakan pada penelitian kali ini adalah *Naive Bayes Algorithm*. Adapun algoritma tersebut digunakan karena berdasarkan penelitian sebelumnya, algoritma *Naive Bayes Algorithm* memiliki akurasi yang tinggi, dan bahkan lebih tinggi apabila dibandingkan dengan *Support Vector Machine (SVM)*. Dan untuk mengoptimalkan model klasifikasi, digunakan teknik *SMOTE (Synthetic Minority Over-sampling Technique)* untuk meningkatkan kinerja model dalam mengatasi ketidakseimbangan kelas atau label pada *dataset* [8].

2. Metodologi Penelitian

Penelitian ini menggunakan *tools Google Colaboratory (Google Colab)* dan menggunakan bahasa pemrograman *Python*. Data yang diperoleh dari proses

scrapping data selanjutnya akan dilakukan analisis sentimen menggunakan *Naive Bayes Algorithm* dengan teknik *SMOTE (Synthetic Minority Over-sampling Technique)* sebagai optimasi model klasifikasinya.

Tahapan penelitian pada penelitian ini diuraikan ke dalam bentuk *flowchart* pada Gambar 1.



Gambar 1. Tahapan Penelitian

2.1 Pengumpulan Data

Pengumpulan data dilakukan melalui proses *Scrapping* Data. Data yang di-*scrapping* yaitu data komentar pengguna dari video di *platform* YouTube Shorts dengan judul “BOIKOT PRODUK ISRAEL satu produk yg kalian beli sama dgn kalian membeli amunisi mereka”. Proses *Scrapping* Data dilakukan dengan bantuan *website* Netlytic, yang hanya membutuhkan *video id* dari video yang akan di-*scrapping* data nya. Data komentar pengguna kemudian di-*export* dalam bentuk format file .CSV (Comma Separated Values).

Data komentar yang berhasil dikumpulkan dari proses *Scrapping* data yaitu sebanyak 2433 data. Namun, data tersebut masih merupakan data *raw* atau data kotor (belum di *pre-processing*). Contoh dari data hasil *scrapping* ditunjukkan pada Tabel 1.

Tabel 1. Contoh data raw (kotor) hasil dari *scrapping*

author	description
@hayuning9612	Produk israel bagus ko 🍌🍌
@Zahwa_28	Itu bukan produk Israel tapi produk Amerika/dukung Israel
@IbuIbu-bn3dc	Facebook wa YouTube Instagram juga buatan yahudi di Amerika Serikat, terus harus diboikot?
@zackygintin g7721	#freepalestineps

2.2 Pelabelan Data

Pelabelan data merupakan proses yang bertujuan untuk menganalisa sifat komentar yang memiliki sifat positif, negatif atau netral [9]. Setelah dilakukan proses *cleaning* pada dataset tahap selanjutnya yaitu melakukan pelabelan data. Pendekatan yang digunakan dalam melakukan *analisis sentimen* terhadap setiap record data ini adalah menggunakan algoritma Bidirectional Encoder Representations from Transformers (BERT) [10]. Pelabelan pada dataset dilakukan bertujuan untuk mempermudah proses training pada data agar mendapatkan hasil akurasi yang baik.

Tabel 2. Hasil pelabelan data menggunakan transformer

description	clean	sentimen
Produk israel	produk israel bagus ko	positive
bagus ko 🍌 🍌		
Tapi tidak semua hal datang dari israel, keban...	tapi tidak semua hal datang dari israel keban...	negative
Nasib kita sama	nasib kita sama	positive
Pertanyaan yg sering dengar...cobalah untuk ...	pertanyaan yg sering dengar cobalah untuk pakai...	neutral

2.3. Data Preprocessing

Data Preprocessing adalah tahap awal dalam pemrosesan teks untuk menyiapkan teks agar dapat diolah lebih lanjut. Sekumpulan karakter yang tersambung (teks) perlu dipecah menjadi elemen-elemen yang lebih bermakna, yang dapat dilakukan pada berbagai tingkatan [11]. Pada tahap ini akan dilakukan *preprocessing* terhadap data komentar sebanyak 2433 baris untuk membersihkan noise yang ada pada data yang meliputi *Case Folding*, *Data Cleaning*, *Tokenization*, *Normalization*, *Stopword Removal* dan *Stemming*.

2.3.1 Case Folding

Case folding adalah tahap mengubah huruf besar menjadi huruf kecil. Proses ini bertujuan untuk mempermudah langkah-langkah pada pemrosesan berikutnya [12]. Hasil pada tahap ini dapat dilihat pada tabel 3.

Tabel 3. Hasil proses case folding

description	Case folding
Menyesal beli 🍌 🍌 🍌	menyesal beli 🍌 🍌 🍌
Semua itu promo dari Israel	semua itu promo dari israel
Emang bango israel?	emang bango israel?
OREO DARI ISRA*L 🍌 🍌 🍌	oreo dari isra*l 🍌 🍌 🍌
PADAHAL ITU KESUKAAN KU LOH 🍌 🍌 🍌	padahal itu kesukaan ku loh 🍌 🍌 🍌

2.3.2 Data Cleaning

Data Cleaning adalah salah satu tahap dari *pre-processing*. Pada tahap ini, data yang telah melalui proses *case folding* akan digunakan untuk menghapus data yang mengandung mention seperti (@username), hashtag (#), URL (http:\\situs.com), karakter non-alfabet, dan emotikon [13]. Hasil pada tahap ini dapat dilihat pada tabel 4.

Tabel 4. Hasil proses data cleaning

Case Folding	Cleaning
menyesal beli 🍌 🍌 🍌	menyesal beli
semua itu promo dari israel	semua itu promo dari israel
emang bango israel?	emang bango israel
oreo dari isra*l 🍌 🍌 🍌	oreo dari isral padahal itu
padahal itu kesukaan ku loh 🍌 🍌 🍌	kesukaan ku loh

2.3.3 Tokenization

Tokenization digunakan untuk menganalisis seluruh kalimat pada data komentar. Selanjutnya, kata-kata dipisahkan berdasarkan karakter pemisahannya, sehingga kata-kata yang bukan karakter pemisah akan digabungkan dengan karakter berikutnya [12]. Hasil pada tahap ini dapat dilihat pada tabel 5.

Tabel 5. Hasil proses tokenization

Case Folding	Cleaning
menyesal beli	['menyesal', 'beli']
semua itu promo dari israel	['semua', 'itu', 'promo', 'dari', 'israel']
emang bango israel	['emang', 'bango', 'israel']
oreo dari isral padahal itu kesukaan ku loh	['oreo', 'dari', 'isral', 'padahal', 'itu', 'kesukaan', 'ku', 'loh']

2.3.4 Normalization

Normalization adalah tahap penting dalam analisis sentimen, terutama ketika menangani data komentar dari platform seperti YouTube Shorts, di mana banyak komentar tidak menggunakan bahasa baku seperti "gw" dan "loe", serta singkatan seperti "yg" dan "gk". Tanpa normalisasi, kata-kata ini akan dianggap sebagai entitas yang berbeda, meskipun memiliki makna yang sama. Oleh karena itu, perlu dilakukan normalisasi untuk mengubah kata-kata tidak baku menjadi bentuk baku [14].

Pada proses penelitian ini digunakan *kamus alay* dari platform *Github* Nikmatul Aliyah Salsabila dalam melakukan normalisasi. Hasil pada tahap ini dapat dilihat pada tabel 6.

Tabel 6. Hasil proses normalization

Tokenization	Normalization
--------------	---------------

['gw', 'gapeduli', 'sih']	['gue', 'gapeduli', 'sih']
['bg', 'aqua', 'kan',	['bang', 'aqua', 'kan',
'produk', 'indo', 'bg',	'produk', 'indonesia',
'bukan', 'boikot']	'bang', 'bukan', 'boikot']

2.3.5 Stopword Removal

Stopword removal merupakan proses penghilangan kata-kata umum yang dianggap tidak berkontribusi pada penentuan sentimen teks [15]. Pada proses penelitian ini digunakan *kamus stopwords* dari platform *Github*. Contoh *Stopword* pada proses ini diantaranya 'Ada', 'adapun', 'agak', 'agaknyanya', 'begini', 'beginian', 'cara', 'caranya', 'cukup', 'cukupkah', 'dahulu', 'dalam', 'dan', 'dapat', 'dari', 'daripada', 'datang', 'dekat', 'entah', 'enggak', 'guna', 'gunakan', 'hal', 'hanya', 'kok'. Hasil dari proses ini dapat dilihat pada tabel 7.

Tabel 7. Hasil proses *stopword removal*

Tokenization	Normalization
['produk', 'israel',	['produk', 'israel',
'bagus', 'kok']	'bagus']
['abc', 'kok', 'enggak',	['abc', 'boykot']
'boykot']	

2.3.6 Stemming

Stemming adalah proses merubah kata -kata berimbuhan menjadi kata dasar [13]. Proses ini bertujuan untuk meningkatkan konsistensi kata, akurasi model, serta menghindari duplikasi kata. Hasil pada tahap ini dapat dilihat pada tabel 8.

Tabel 8. Hasil proses *stemming*

Normalization	Stemming
['mengatur', 'hak',	['atur', 'hak', 'pilih']
'memilih']	
['memblokir', 'produk',	['blokir', 'produk',
'mendukung', 'israel']	'dukung', 'israel']

2.5. Wordcloud

Wordcloud (juga dikenal sebagai tag cloud) dikenali sebagai gambar yang terdiri dari kata-kata yang digunakan dalam teks atau topik tertentu di mana frekuensi atau maknanya ditunjukkan oleh ukuran setiap kata [16].

Wordcloud dibuat menggunakan library 'wordcloud' pada bahasa pemrograman *python*. Dalam penelitian ini, *wordcloud* digunakan untuk melakukan visualisasi pada data yang telah di-*preprocessing*.

2.4. Pembobotan Kata

Pembobotan kata dilakukan dengan menggunakan teknik TF-IDF. TF-IDF(Term Frequency-Inverse Document Frequency) merupakan suatu metode algoritma yang memberikan bobot terhadap teks [15]. Pada penelitian ini library *python TF-IDF Vectorizer* digunakan untuk melakukan proses pembobotan kata.

TF adalah frekuensi suatu kata muncul dalam dokumen, sedangkan IDF adalah nilai invers dari dokumen yang mengandung kata tersebut. TF dan IDF akan dikalikan sehingga menghasilkan nilai bobot dari kata tersebut [17]. Berikut merupakan persamaan TF-IDF :

$$TF - IDF(d, t) = TD(d, t) \times IDF(t) \quad (1)$$

dimana :

$$TF(d, t) = \frac{\text{jumlah kata } t \text{ pada dokumen } d}{\text{total kata pada dokumen } d} \quad (2)$$

$$IDF(t) = \frac{\text{total dokumen}}{\text{jumlah dokumen yang mengandung kata } t} \quad (3)$$

Keterangan :

t = kata

d = dokumen

2.5 SMOTE

SMOTE (Synthetic Minority Over-sampling Technique) digunakan untuk mengatasi masalah ketidakseimbangan kelas dengan mensintesis sampel baru dari kelas minoritas. Teknik ini membuat sampel baru dari kelas minoritas dengan membentuk kombinasi konveks dari instance tetangga. Tingkat oversampling yang diinginkan kemudian dipilih secara acak. Teknik ini menggambar garis antara titik-titik minoritas di ruang fitur dan menambahkan sampel di sepanjang garis tersebut. Contoh baru dari kelas minoritas ini ditambahkan ke data training, dan klasifikator dilatih dengan data tambahan. Algoritma SMOTE secara umum lebih akurat dibandingkan dengan metode *oversampling* biasa [18].

2.6 Naive Bayes

Naïve Bayes classifier merupakan suatu metode klasifikasi yang menggunakan perhitungan probabilitas. Konsep dasar yang digunakan pada *Naïve Bayes* classifier adalah Teorema Bayes yang dinyatakan pertama kali oleh Thomas Bayes [19]. Algoritma ini digunakan dalam melakukan implementasi dalam penelitian ini, dengan berdasarkan pada pengelompokan *data training* dan *data uji*. Nilai probabilitas yang digunakan dinyatakan secara sederhana sebagai berikut :

$$P(C | X) = \frac{P(X|C).P(C)}{P(X)} \quad (4)$$

2.7 Confusion Matrix

Confusion Matrix (Matriks Kebingungan) digunakan untuk mengevaluasi performa dari metode/model klasifikasi yang digunakan. *Confusion Matrix* terdiri dari : *True Positive (TP)*, *True Negative (TN)*, *False Positive (FP)*, dan *False Negative (FN)* yang merupakan representasi dari hasil klasifikasi [20] .

11

sentimen positif dan negatif dalam teks yang dianalisis. Semakin tinggi bobot suatu kata dalam TF-IDF, semakin besar kontribusinya terhadap kemampuan model Naive Bayes yang sudah dilatih dalam mengenali dan mengklasifikasikan sentimen secara akurat.

Sample dari hasil ekstraksi fitur menggunakan TF-IDF pada pembagian data 60% *data training* dan 40% *data testing* dapat dilihat pada Tabel 9.

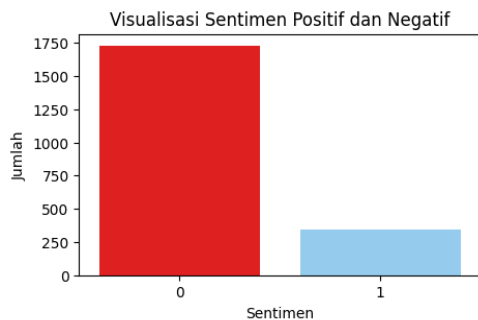
Tabel 9. Hasil ekstraksi fitur dengan teknik TF-IDF

Dokumen	Kata	Nilai TF-IDF
0	bango	0.632696
0	boikot	0.283752
0	indonesia	0.354839
0	kecap	0.627113
1	beli	0.101252
...	...	...
830	israel	0.225329
830	konten	0.228782
830	lindung	0.357278
830	orang	0.297105
830	tiada	0.357278
...	...	...

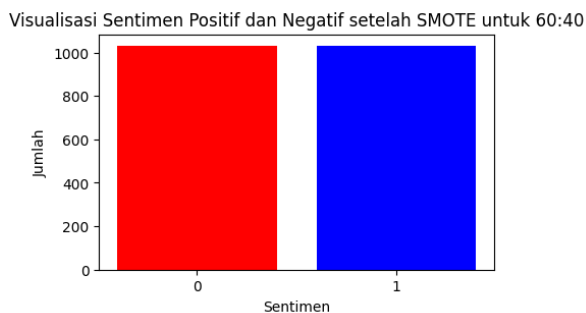
3.1.3 Hasil Penyeimbangan Data

Penyeimbangan data pada penelitian ini dilakukan menggunakan teknik SMOTE (*Synthetic Minority Over-sampling Technique*). Teknik SMOTE digunakan pada penelitian ini untuk menyeimbangkan data yang digunakan pada analisis sentimen sehingga model klasifikasi menjadi lebih optimal.

Visualisasi grafik dari data sebelum dan sesudah dioptimalkan dengan Teknik SMOTE pada pembagian data 60% *data training* dan 40% *data testing* dapat dilihat pada Gambar 4 dan Gambar 5.



Gambar 4. Visualisasi Data sebelum penyeimbangan

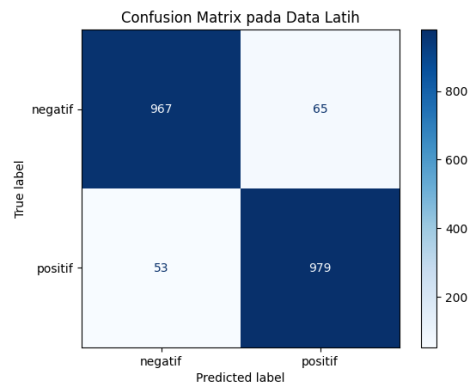


Gambar 5. Visualisasi Data setelah penyeimbangan dengan Teknik SMOTE

3.1.4 Hasil Confusion Matrix

Hasil dari *Confusion Matrix* memberikan gambaran tentang kinerja model *Naive Bayes* dalam mengklasifikasikan sentimen pada data komentar. *Confusion Matrix* membagi prediksi ke dalam empat kategori: True Positives (TP), di mana sentimen positif secara akurat diidentifikasi sebagai positif; True Negatives (TN), di mana sentimen negatif secara benar diidentifikasi sebagai negatif; False Positives (FP), kesalahan dimana sentimen negatif salah diklasifikasikan sebagai positif; dan False Negatives (FN), di mana sentimen positif tidak terdeteksi dan salah diklasifikasikan sebagai negatif. *Confusion Matrix* pada penelitian ini membantu dalam mengukur secara tepat berapa banyak kesalahan yang dibuat oleh model serta keberhasilannya dalam mengidentifikasi sentimen pada data komentar.

Sample dari *Confusion Matrix* pada pembagian data 60% *data training* dan 40% *data testing* dapat dilihat pada Gambar 6.



Gambar 6. *Confusion Matrix* dari kinerja model pada data 60:40

3.1.4 Hasil Implementasi Naive Bayes Algorithm

Hasil dari implementasi *Naive Bayes Algorithm* dalam penelitian ini yaitu dikelompokkan berdasarkan pembagian data, yaitu dengan proporsi 80% *data training* dan 20% *data testing*, 60% *data training* dan 40% *data testing*, 90% *data training* dan 10% *data testing* dan dibedakan menjadi dua jenis, yaitu dengan menggunakan data yang sudah di-seimbangkan dengan teknik SMOTE dan yang tidak.

Akurasi *Naive Bayes Algorithm* pada proporsi pembagian data 80% *data training* dan 20% *data testing* pada data yang sudah diseimbangkan dengan teknik SMOTE mencapai tingkat *Accuracy* (akurasi) sebesar 0.9325309992706053 atau 93%. Sedangkan untuk data yang tidak melalui proses penyeimbangan pada proporsi pembagian data yang sama mencapai tingkat *Accuracy* (akurasi) sebesar 0.8356411800120409 atau 84%.

Tabel 10 menunjukkan hasil pengujian pada data yang telah diseimbangkan dengan teknik SMOTE dan Tabel 11 menunjukkan hasil pengujian pada data yang tidak

melalui proses penyeimbangan pada proporsi data 80% *data training* dan 20% *data testing*.

Tabel 10. Hasil Implementasi Pada Data 80:20 dengan teknik SMOTE

	precision	recall	f1-score
negative	0.92	0.95	0.93
positive	0.95	0.91	0.93
accuracy			0.93
macro avg	0.93	0.93	0.93
weighted avg	0.93	0.93	0.93

Tabel 11. Hasil Implementasi Pada Data 80:20 tanpa proses penyeimbangan

	precision	recall	f1-score
negative	0.83	1.00	0.91
positive	1.00	0.06	0.11
accuracy			0.94
macro avg	0.92	0.53	0.51
weighted avg	0.86	0.84	0.77

Akurasi *Naive Bayes Algorithm* pada proporsi pembagian data 60% *data training* dan 40% *data testing* pada data yang sudah diseimbangkan dengan teknik SMOTE mencapai tingkat *Accuracy* (akurasi) sebesar 0.9428294573643411 atau 94%. Sedangkan untuk data yang tidak melalui proses penyeimbangan pada proporsi pembagian data yang sama mencapai tingkat *Accuracy* (akurasi) sebesar 0.8378812199036918 atau 84%.

Tabel 12 menunjukkan hasil pengujian pada data yang telah diseimbangkan dengan teknik SMOTE dan Tabel 13 menunjukkan hasil pengujian pada data yang tidak melalui proses penyeimbangan pada proporsi data 60% *data training* dan 40% *data testing*.

Tabel 12. Hasil Implementasi Pada Data 60:40 dengan teknik SMOTE

	precision	recall	f1-score
negative	0.95	0.94	0.94
positive	0.94	0.95	0.94
accuracy			0.94
macro avg	0.94	0.94	0.94
weighted avg	0.94	0.94	0.94

Tabel 13. Hasil Implementasi Pada Data 60:40 tanpa proses penyeimbangan

	precision	recall	f1-score
negative	0.84	1.00	0.91
positive	1.00	0.06	0.11
accuracy			0.84
macro avg	0.92	0.53	0.51
weighted avg	0.86	0.84	0.77

Akurasi *Naive Bayes Algorithm* pada proporsi pembagian data 90% *data training* dan 10% *data testing* pada data yang sudah diseimbangkan dengan teknik SMOTE mencapai tingkat *Accuracy* (akurasi) sebesar 0.9368556701030928 atau 94%. Sedangkan untuk data yang tidak melalui proses penyeimbangan pada proporsi pembagian data yang sama mencapai

tingkat *Accuracy* (akurasi) sebesar 0.8400214018191546 atau 84%.

Tabel 14 menunjukkan hasil pengujian pada data yang telah diseimbangkan dengan teknik SMOTE dan Tabel 15 menunjukkan hasil pengujian pada data yang tidak melalui proses penyeimbangan pada proporsi data 90% *data training* dan 10% *data testing*.

Tabel 14. Hasil Implementasi Pada Data 90:10 dengan teknik SMOTE

	precision	recall	f1-score
negative	0.92	0.95	0.94
positive	0.95	0.92	0.94
accuracy			0.94
macro avg	0.94	0.94	0.94
weighted avg	0.94	0.94	0.94

Tabel 15. Hasil Implementasi Pada Data 90:10 tanpa proses penyeimbangan

	precision	recall	f1-score
negative	0.84	1.00	0.91
positive	1.00	0.06	0.11
accuracy			0.84
macro avg	0.92	0.53	0.51
weighted avg	0.87	0.84	0.78

3.2 Pembahasan

Penelitian analisis sentimen ini dilakukan dengan menggunakan algoritma *Naive Bayes* sebagai model untuk memprediksi apakah teks pada komentar cenderung memiliki sentimen positif atau negatif. Adapun algoritma *Naive Bayes* pada penelitian ini dilatih dengan *dataset* komentar publik sebanyak 2433 data komentar yang diambil dari *platform* berbagi video YouTube Shorts.

Setelah data komentar diambil, dilakukan tahap *pre-processing* pada data komentar dan didapat bahwa sentimen masyarakat Indonesia terhadap aksi boikot produk Israel cenderung negatif, seperti yang dapat dilihat pada Gambar 4 dimana data komentar yang dikategorikan sebagai komentar negatif jauh lebih banyak apabila dibandingkan dengan data komentar yang dikategorikan sebagai komentar positif.

Selanjutnya, Model *Naive Bayes* dilatih menggunakan data komentar yang telah melalui proses *pre-processing*. Data tersebut kemudian dibagi menjadi beberapa proporsi untuk *training* dan *testing*, yaitu 80% *data training* dan 20% *data testing*, 60% *data training* dan 40% *data testing*, serta 90% *data training* dan 10% *data testing*. Pengujian model ini juga dilakukan dalam dua kondisi berbeda, yaitu menggunakan data yang telah diseimbangkan dengan teknik SMOTE dan data yang tidak diseimbangkan.

Hasil pengujian menunjukkan bahwa model yang dilatih dengan data yang diseimbangkan menggunakan teknik SMOTE pada proporsi 60% *data training* dan 40% *data testing* mencapai akurasi tertinggi dalam analisis sentimen teks komentar di penelitian ini, yaitu sebesar 94% atau 0.9428294573643411. Sebaliknya,

ketika model diuji dengan proporsi pembagian data yang sama tetapi pada data yang tidak diseimbangkan, tingkat akurasi model turun menjadi 84% atau 0.8378812199036918.

Sedangkan, pengujian dengan proporsi 90% *data training* dan 10% *data testing* pada data yang tidak diseimbangkan, tingkat akurasi model mencapai 0.8400214018191546 atau 84%, sedikit lebih tinggi apabila dibandingkan dengan pengujian yang sama pada proporsi 60% *data testing* dan 40% *data training*. Tetapi ketika model diuji dengan proporsi pembagian data yang sama namun data diseimbangkan terlebih dahulu menggunakan teknik SMOTE, tingkat akurasi model meningkat yaitu mencapai 0.9368556701030928 atau 94%, walaupun sedikit rendah apabila dibandingkan dengan pengujian yang sama pada proporsi 60% *data testing* dan 40% *data training*.

4. Kesimpulan

Berdasarkan hasil dari penelitian yang sudah dilakukan, penelitian ini menunjukkan bahwa Algoritma *Naive Bayes* memiliki keefektifan yang tinggi dalam memprediksi sentimen dari suatu teks, dengan tingkat akurasi mencapai 94%. Akan tetapi, penggunaan teknik penyeimbangan data yaitu teknik SMOTE sangat mempengaruhi tingkat akurasi dari model *Naive Bayes* dalam analisis sentimen pada suatu kalimat/kata.

Penggunaan SMOTE terbukti meningkatkan kemampuan model dalam mengklasifikasikan sentimen dengan lebih akurat. Dengan menyediakan data yang lebih seimbang, teknik SMOTE membantu mengatasi bias yang mungkin muncul akibat ketidakseimbangan kelas yang signifikan, sehingga meningkatkan keefektifan model dalam mengklasifikasi sentimen pada suatu kalimat/kata.

Penelitian ini juga mengungguli penelitian sebelumnya, dikarenakan lebih banyaknya jumlah *dataset* yang digunakan serta penggunaan algoritma *BERT* untuk pelabelan data dan teknik *SMOTE* untuk penyeimbangan data, yang secara signifikan meningkatkan akurasi prediksi sentimen pada model yang dilatih.

Dan, dari penelitian ini juga terungkap bahwa sentimen masyarakat Indonesia terhadap aksi boikot produk Israel cenderung negatif.

Daftar Pustaka

- [1] O.: Mohd, R. Mohd, and N. □ Abstrak, "KONFLIK ISRAEL-PALESTIN DARI ASPEK SEJARAH MODEN DAN LANGKAH PEMBEBASAN DARI CENGKAMAN ZIONIS."
- [2] M. F. Millenio, "How the Judgement Effective? The Role of United Nations in Conflict Resolution Between Palestine and Israel," *The Digest: Journal of Jurisprudence and*
- [3] *Legisprudence*, vol. 2, no. 2, pp. 197–230, 2021.
- [4] M. Wendra and A. Sutrisno, "Tantangan Penyelesaian Konflik Internasional yang Dilematik mengenai Hak Veto dalam Dewan Keamanan PBB (Studi kasus Palestina dengan Israel) ARTICLE HISTORY," *Journal of Contemporary Law Studies*, no. 2, pp. 171–180, 2024, doi: 10.47134/lawstudies.v2i2.2373.
- [5] B. Liu, "Sentiment Analysis and Subjectivity."
- [6] A. T. Susilawati, N. A. Lestari, and P. A. Nina, "Analisis Sentimen Publik Pada Twitter Terhadap Boikot Produk Israel Menggunakan Metode Naïve Bayes," *Nian Tana Sikka: Jurnal Ilmiah Mahasiswa*, vol. 2, no. 1, pp. 26–35, 2024.
- [7] N. F. Az-haari, D. Juardi, and A. Jamaludin, "Analisis Sentimen Terhadap Boikot Brand Pro-Israel pada Twitter Menggunakan Metode Naïve Bayes," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 3, pp. 4256–4261, 2024.
- [8] L. A. Hayurian and N. Hendrastuty, "COMPARISON OF NAÏVE BAYES ALGORITHM AND SUPPORT VECTOR MACHINE IN SENTIMENT ANALYSIS OF BOYCOTT ISRAELI PRODUCTS ON TWITTER," vol. 5, no. 3, pp. 731–738, 2024, doi: 10.52436/1.jutif.2024.5.3.1813.
- [9] M. N. Huda, D. A. Fauzan, M. R. S. P. Pamungkas, N. S. Ratnadewi, and A. A. Vahendra, "Optimalisasi Model Klasifikasi Sentimen Netizen Terhadap Merek Tas Luar Negeri," *Jurnal KomtekInfo*, pp. 21–28, Mar. 2023, doi: 10.35134/komtekinfo.v10i1.360.
- [10] E. Kontopoulos, C. Berberidis, T. Dergiades, and N. Bassiliades, "Ontology-based sentiment analysis of twitter posts," *Expert Syst Appl*, vol. 40, no. 10, pp. 4065–4074, Aug. 2013, doi: 10.1016/j.eswa.2013.01.001.
- [11] C. Alifia Putri and S. Al Faraby, "Analisis Sentimen Review Film Berbahasa Inggris Dengan Pendekatan Bidirectional Encoder Representations from Transformers," vol. 6, no. 2, pp. 181–193, 2020, [Online]. Available: <http://jurnal.mdp.ac.id>
- [12] R. Noviana and I. Rasal B A Jurusan, "PENERAPAN ALGORITMA NAIVE BAYES DAN SVM UNTUK ANALISIS SENTIMEN BOY BAND BTS PADA MEDIA SOSIAL TWITTER," *JTS*, vol. 2, no. 2.
- [13] A. C. Khotimah and E. Utami, "COMPARISON NAÏVE BAYES CLASSIFIER, K-NEAREST NEIGHBOR AND SUPPORT VECTOR MACHINE IN THE CLASSIFICATION OF INDIVIDUAL ON TWITTER ACCOUNT," *Jurnal Teknik Informatika (JUTIF)*, vol. 3, no. 3, 2022, doi: 10.20884/1.jutif.2022.3.3.254.

- [13] N. Hendrastuty, A. Rahman Isnain, and A. Yanti Rahmadhani, "Analisis Sentimen Masyarakat Terhadap Program Kartu Prakerja Pada Twitter Dengan Metode Support Vector Machine," vol. 6, no. 3, 2021, [Online]. Available: <http://situs.com>
- [14] N. Saputra, T. B. Adji, and A. E. Permanasari, "ANALISIS SENTIMEN DATA PRESIDEN JOKOWI DENGAN PREPROCESSING NORMALISASI DAN STEMMING MENGGUNAKAN METODE NAIVE BAYES DAN SVM Oleh," 2015.
- [15] M. I. Fikri, T. S. Sabrila, and Y. Azhar, "Perbandingan Metode Naïve Bayes dan Support Vector Machine pada Analisis Sentimen Twitter," *SMATIKA JURNAL*, vol. 10, no. 02, pp. 71–76, Dec. 2020, doi: 10.32664/smatika.v10i02.455.
- [16] V. Ibrahim, J. A. Bakar, N. H. Harun, and A. F. Abdulateef, "A word cloud model based on hate speech in an online social media environment," *Baghdad Science Journal*, vol. 18, pp. 937–946, Jun. 2021, doi: 10.21123/bsj.2021.18.2(Suppl.).0937.
- [17] B. K. Hananto, A. Pinandito, and A. P. Kharisma, "Penerapan Maximum TF-IDF Normalization Terhadap Metode KNN Untuk Klasifikasi Dataset Multiclass Panichella Pada Review Aplikasi Mobile," 2018. [Online]. Available: <http://j-ptiik.ub.ac.id>
- [18] Hermanto, A. Y. Kuntoro, T. Asra, E. B. Pratama, L. Effendi, and R. Ocanitra, "Gojek and Grab User Sentiment Analysis on Google Play Using Naive Bayes Algorithm and Support Vector Machine Based Smote Technique," in *Journal of Physics: Conference Series*, IOP Publishing Ltd, Nov. 2020. doi: 10.1088/1742-6596/1641/1/012102.
- [19] A. Z. Amrullah, A. Sofyan Anas, M. Adrian, and J. Hidayat, "Analisis Sentimen Movie Review Menggunakan Naive Bayes Classifier Dengan Seleksi Fitur Chi Square," *Jurnal*, vol. 2, no. 1, 2020, doi: 10.30812/bite.v2i1.804.
- [20] H. Yun, "Prediction model of algal blooms using logistic regression and confusion matrix," *International Journal of Electrical and Computer Engineering*, vol. 11, no. 3, pp. 2407–2413, Jun. 2021, doi: 10.11591/ijece.v11i3.pp2407-2413.
- [21] R. Merdiansah and A. Ali Ridha, "Analisis Sentimen Pengguna X Indonesia Terkait Kendaraan Listrik Menggunakan IndoBERT," *Jurnal Ilmu Komputer dan Sistem Informasi (JIKOMSI)*, vol. 7, no. 1, pp. 221–228, 2024.
- [22] D. Valero-Carreras, J. Alcaraz, and M. Landete, "Comparing two SVM models through different metrics based on the confusion matrix," *Comput Oper Res*, vol. 152, Apr. 2023, doi: 10.1016/j.cor.2022.106131.