

PENERAPAN ALGORITMA NAÏVE BAYES DAN K-NEAREST NEIGHBOR UNTUK ANALISIS SENTIMEN YOUTUBE MENGENAI INTENSIF MOBIL LISTRIK

Caswadi ¹, Nisa Dienwati ², Gifthera Dwilestari ³, Fathurrohman ⁴, Edi Tohidi ⁵

^{1,2} Teknik Informatika, STMIK IKMI Cirebon

³ Sistem Informasi, STMIK IKMI Cirebon

⁴ Rekayasa Perangkat Lunak, STMIK IKMI Cirebon

⁵ Komputerisasi Akuntansi, STMIK IKMI Cirebon

Jalan Perjuangan 10B, Karyamulya, Kec. Kesambi, Kota Cirebon, Jawa Barat

caswadip@gmail.com

ABSTRAK

Dampak pemanasan global terhadap perubahan iklim telah menarik perhatian dari berbagai pihak. 27% dari polusi udara disebabkan oleh transportasi. Di berbagai negara, pemerintah telah menerapkan langkah-langkah untuk mengurangi polusi udara dengan mendorong masyarakat untuk beralih menggunakan kendaraan listrik. Namun keberhasilan pemerintah untuk mengkampanyekan teknologi mobil listrik tergantung pada sentimen dan pemahaman masyarakat. Tujuan penelitian ini untuk menganalisa sentimen publik. Dataset digunakan dalam penelitian ini ialah 1517 data komentar di laman *youtube* mengenai intensif mobil listrik. Berdasarkan analisis, sebagian besar komentar *youtube* memiliki sentimen negatif (57,4%), sementara jumlah komentar yang positif (33,3%) dan netral (9,3%). Penelitian ini menerapkan metode *Naïve Bayes* dan *K-Nearest Neighbor*. Hasil penelitian menunjukkan bahwa metode *K-Nearest Neighbor* memberikan hasil klasifikasi yang lebih baik dengan *Accuracy* sebesar 93,23%, *precision* 93,91% dan *recall* 91,56%. Sedangkan *Naïve Bayes* memperoleh hasil *Accuracy* 86,95%, *precision* 80,51%, dan *recall* 91,23%.

Kata kunci: Mobil Listrik, Klasifikasi, Analisis Sentimen, Youtube.

1. PENDAHULUAN

Dalam upaya untuk mengurangi emisi gas di negara-negara anggota G20, dengan mengacu pada target *Nasional Determined Contribution (NDC)*. *Climate Transparency*, pada laporan tahunan G20 mengenai transparansi iklim 2022 [1] Berdasarkan data tersebut, dapat disimpulkan bahwa transportasi merupakan penyumbang sebanyak 27 persen terhadap polusi udara dunia. Kontributor utama yang signifikan terhadap pencemaran udara adalah transportasi yang menggunakan bahan bakar fosil, termasuk penggunaan alat berat. Sarana yang dapat diadopsi guna mengurangi dan meredam emisi gas melibatkan implementasi kendaraan listrik (*electric vehicle*), suatu pendekatan yang telah diupayakan secara intensif oleh beberapa negara belakangan ini. Pemerintah Indonesia mengimplementasikan regulasi terkait penggunaan kendaraan listrik dalam konteks transportasi jalan, yang mencakup kendaraan berbaterai. Langkah ini dianggap sebagai bukti konkret dari komitmen pemerintah untuk meningkatkan upaya dalam menekan emisi gas dan mengatasi dampak perubahan iklim.

Pertama kali, mobil listrik diproduksi secara massal dan komersial oleh *General Motors* dirancang untuk digunakan pada durasi waktu yang terbatas dan menempuh jarak yang relatif pendek bila dibandingkan dengan kendaraan konvensional [2]. Tidak hanya terkait dengan aspek teknis, tetapi juga perlu memperhatikan persepsi sosial dalam rangka meningkatkan distribusi komersial kendaraan listrik. Faktor yang mempengaruhi tingkat adopsi mobil listrik di Indonesia mencakup persepsi masyarakat

terhadap kendaraan listrik. Adanya korelasi positif, antara perasaan positif saat mengendarai kendaraan listrik dan niat konsumen. Sementara di sisi lain, beberapa konsumen melaporkan pengalaman perasaan negatif seperti “merasa lebih boros” setelah menggunakan kendaraan listrik [3].

Berdasarkan penelitian yang pernah dilakukan oleh Riyadi dkk dengan topik analisa pengukuran sentimen *twitter* mengenai teknologi kendaraan listrik diketahui metode *Naïve Bayes*, *K-Nearest Neighbor* dan *Decision Tree* menghasilkan tingkat *Accuracy* 94% dengan respon positif 53%, negatif sebesar 38% dan nilai respon netral 9%, yang artinya kendaraan mobil listrik di Indonesia tetap dapat diterima oleh masyarakat [3]. Penelitian terkait lainnya yaitu Mengimplementasikan algoritma *Naïve Bayes*, *Support Vector Machine*, dan *K-Nearest Neighbor* untuk menganalisis sentimen aplikasi Halodoc [4]. Data yang diambil menggunakan *scraping* pada *Google Play Store*, data yang berhasil terkumpul adalah 500 data dengan ulasan positif dan 500 data dengan ulasan negatif. Hasil penelitian menunjukkan algoritma *K-Nearest Neighbor* adalah algoritma terbaik dalam penelitian tersebut dengan tingkat *accuracy* 95.0%.

Penelitian ini berfokus pada analisis sentimen terhadap data komentar *youtube* menggunakan algoritma *Naïve Bayes* dan *K-Nearest Neighbor*. Metode penelitian ini akan mencakup pengumpulan data komentar *youtube* yang relevan dengan topik, *preprocessing* data untuk membersihkan dan mempersiapkan teks, ekstraksi fitur dari teks, dan penggunaan algoritme *Naïve Bayes* dan *K-Nearest*

Neighbor untuk melakukan klasifikasi sentimen. Algoritma *Naïve Bayes* digunakan untuk mengklasifikasikan teks menjadi kategori sentimen berdasarkan probabilitas, sementara *K-Nearest Neighbor* digunakan untuk mengelompokkan komentar youtube ke dalam klaster berdasarkan kesamaan mereka.

2. TINJAUAN PUSTAKA

2.1. Text Mining

Text mining adalah suatu proses ekstraksi pola-pola yang signifikan dan menarik dari data tekstual sumber dengan tujuan untuk mengeksplorasi pengetahuan. *Text mining* merupakan bidang multidisiplin yang mendasarkan diri pada pengambilan informasi, penambahan data, pembelajaran mesin, statistik, dan linguistik komputasi. *Text mining* juga dapat dipahami sebagai upaya pencarian data teks yang melibatkan keinginan untuk menemukan pengetahuan yang sebelumnya belum teridentifikasi. Proses ini merupakan elemen esensial dari pencarian data teks yang bertujuan untuk memperoleh pemahaman, mengidentifikasi pola praktis dan merumuskan pengetahuan dari kumpulan data teks atau corpus masif (kumpulan teks yang menangkap penggunaan bahasa dalam bentuk tertulis atau lisan secara utuh dan padat) yang bersifat tidak terstruktur [6].

2.2. Naïve Bayes

Naïve Bayes merupakan salah satu metode yang dapat diterapkan untuk menganalisis sentimen, dan secara teoritis, metode ini menunjukkan konsistensi data dan ketepatan perhitungan klasifikasi. Penggunaan *Naïve Bayes* umumnya terfokus pada teknik klasifikasi, terutama dalam konteks laman seperti *twitter*. Ciri khas dari klasifikasi *Naïve Bayes* terletak pada kemampuannya untuk menghasilkan hipotesis yang kuat terkait suatu kondisi atau kejadian. Perhitungan probabilitas kelompok dalam metode *Naïve Bayes* mengacu pada pendekatan algoritma *Bayesian* yang menggunakan persamaan sebagai dasarnya [7].

2.3. K-Nearest Neighbor

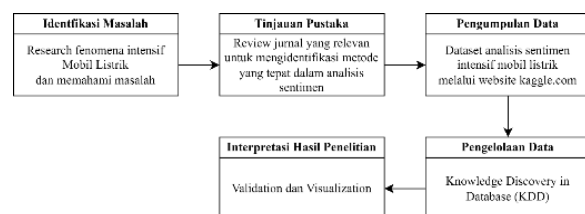
K-Nearest Neighbor adalah algoritma pembelajaran mesin terawasi yang sebagian besar digunakan untuk tujuan klasifikasi. Umumnya algoritma *K-NN* mampu mengklasifikasikan dataset menggunakan model pelatihan yang mirip dengan kueri pengujian dengan mempertimbangkan *K-Nearest* data pelatihan terdekat (*Neighbor*) yang paling dekat dengan kueri yang diuji. algoritma *K-NN* merupakan salah satu bentuk yang paling sederhana dan banyak digunakan dalam tugas-tugas klasifikasi karena memiliki desain yang sangat adaptif dan mudah dipahami [8].

2.4. Knowledge Discovery in Database (KDD)

KDD merupakan kegiatan interdisiplin yang merujuk pada metodologi untuk memperoleh pengetahuan dari data [9]. Penambahan data dan KDD sering digunakan secara bergantian untuk menggambarkan proses penggalian informasi tersembunyi dalam database yang luas [10].

3. METODE PENELITIAN

Metode penelitian ini menggunakan metode penelitian eksperimen kuantitatif. Metode eksperimen adalah metode penelitian ilmiah yang dilakukan dengan merancang situasi atau kondisi tertentu untuk mengukur pengaruh variabel tertentu terhadap variabel lain. Dalam metode penelitian eksperimen dapat menguji sebab akibat secara langsung dan dapat mengontrol variabel yang mempengaruhi hasil penelitian. Dapat dilakukan dengan menggunakan algoritma *Naïve Bayes* dan *K-Nearest Neighbor*.



Gambar 1. Penelitian

4. HASIL DAN PEMBAHASAN

4.1. Hasil Penelitian

4.1.1. Data Selection

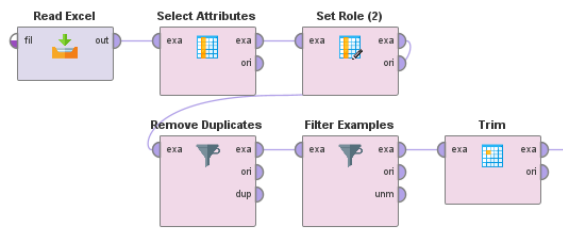
Dataset pada penelitian ini didapat dari website *Kaggle.com* dengan operator search *twitter* dan kata kunci Analisis Sentimen Mobil Listrik. Pengambilan data ini dilakukan pada 1 Desember 2023. Data yang didapat sebanyak 1517 data *tweets*. Bentuk dataset dapat dilihat pada gambar 1 berikut.

Row No.	id Konten	Nama Akun	Tanggal	Teks	Sentimen
1	Ug9B5eF5C3gKJ4k4n4g	Spn Lb	2023-08-05 12:54:49+00:00	siaran sbi bkn harga onst sama kapi bno insu abah lara manis	positif
2	Ug8EUA5C3T5n4C3p4n4g	kuhan ara	2023-08-04 12:10:23+00:00	problem subside kualitas dibarengi harga dinamis usaha gbu dan oan subside sebat off.	negatif
3	Ug9p4u5MF4E4C3Cv4n4g	Falkh Al-Azidi	2023-08-04 10:17:57+00:00	baik kualitas hambang duu baik kualitas motor motor pabikan jepang	positif
4	Ug9VCC8B5n4w4C3T4n4g	sp office	2023-08-04 08:28:54+00:00	model jeni kualitas baik harga mahal onst	negatif
5	Ug9k4u4w4C3Q4n4g	Lembar kuring	2023-08-04 07:55:37+00:00	suarat ngaco woi anak muda boro punya rumah boro jd urum boro boro koi dapat ng	negatif
6	Ug9m4w4E4J4B4n4g	Suati kerangka	2023-08-04 06:58:17+00:00	harga motor mahal Hana harga meng mado keat kualitas bagus banget	positif
7	Ug9m4p4u4C3H4n4g	Sigatun	2023-08-04 06:31:56+00:00	moi karan sbi bkn banyk putnah mo! boro mo! boro boro boro boro boro boro boro	negatif
8	Ug9m4D4C3C4n4g	Pakal Peralda	2023-08-04 01:04:18+00:00	prosesi neral produk banyk subit wala gant kendera boro jdi kendera subit wala boro	negatif
9	Ug9k4C4u4L4F4n4g	Hana Praxido	2023-08-03 23:56:55+00:00	subside kapal sahan	netral
10	Ug9C3T4C3T4n4g	jenan nio asa kate	2023-08-03 11:25:57+00:00	adit rata kerna subside jangen paka gpa mah subside pemilih kherat pemilih	negatif
11	Ug9B5eF5C3gKJ4k4n4g	Randy Ramadnan	2023-08-03 07:20:29+00:00	kapal katar doris kapal belum kapal subside hang boro jd waga hang mampai pila	negatif
12	Ug9m4w4E4J4B4n4g	kuhan ara	2023-08-03 05:08:13+00:00	mungkin subside baik arah engang sebat bno lara gura engang sebat hapu ne	negatif
13	Ug9m4p4u4C3H4n4g	gema	2023-08-03 04:23:25+00:00	kampung sekatang banyk banget boro kono sama cewe pake motor	positif
14	Ug9p4u5MF4E4C3Cv4n4g	F	2023-08-01 12:38:54+00:00	banyk kenderaan boro boro boro boro boro boro boro boro boro boro boro boro boro	negatif
15	Ug9p4u5MF4E4C3Cv4n4g	Khuwidi 22	2023-07-28 11:08:07+00:00	harga kaku mahal	positif

Gambar 2. Dataset

4.1.2. Preprocessing

Untuk meminimalisir adanya *missing value* atau nilai kosong dan data yang tidak konsisten. Pada tahap ini data yang sudah siap diolah akan dihilangkan tanda baca, URL dan *noisy*. Berikut tahapan-tahapan *preprocessing* yaitu menggunakan operator sebagai berikut :

Gambar 3. Proses *Preprocessing*

Adapun penjelasan dari masing-masing operator yaitu:

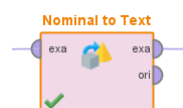
1. *Read Excel*
Operator *Read Excel* digunakan untuk membuka dan membaca dataset dalam bentuk *file excel*
2. *Selection Atribut*
Operator selection berfungsi untuk menyeleksi atribut yang akan digunakan dan menghilangkan atribut lain yang tidak digunakan
3. *Set Role*
Operator *Set Role* digunakan untuk memilih kolom atribut yang akan digunakan sebagai data label
4. *Remove Duplicates*
Operator ini digunakan untuk menghilangkan atau menghapus data yang dinilai memiliki nilai yang sama
5. *Filter Example*
Operator ini digunakan untuk menghilangkan atau menghapus data yang mengandung *missing value* atau data yang tidak memiliki nilai.
6. *Trim*
Operator ini digunakan untuk menghilangkan spasi awal dan akhir dari nilai atribut nominal yang dipilih dan menghilangkan atribut lain yang tidak digunakan.

Berikut hasil dari preprocessing data pada gambar

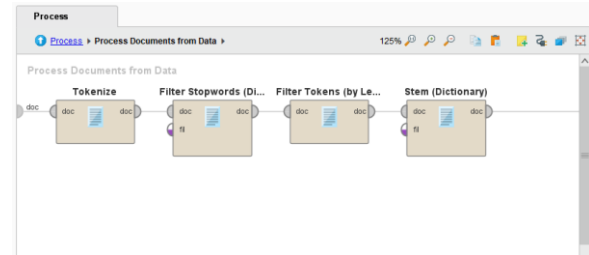
Row No.	Sentimen	Teks
1	positif	caran sih bikin harga nasi sama kayak itu masa sih bare manis
2	negatif	problem outdoor kualitas dikurusi harga dipotong ushah gitu cari ruan outdoor sebis inflasi paling gede
3	positif	bah kualitas nembang dulu baik kualitas motor motor pabrikan jorag
4	negatif	model jeni keawalan banal harga mahal crost
5	negatif	sekarang woi anak muda blom punya rumah blom jd urum buikan seba kar dapat ngapoi ora deui ora deui nu nama outdoor leh an ang maha
6	positif	harga motor mahal masa harga ming motor bakal kualitas bagal tangki bumi
7	negatif	mai nenen jeh bento plus padahal mai lokal ment benera jua pansi tahun anggap bade awet tahun ralat indomesta per tahun sekali kelua jua buat beli bade baru o caran uang
8	negatif	proses rental mobil baru bade waktu gara kendara bon jadi kendara budu waktu teknolog baru sadar
9	netral	suburid head case
10	negatif	adli rata semua outdoor jengon anggar paku paku mod outdoor penatlah kharal pancerda nita lima adli seluruh ralat indomesta ming mau rala bagus ora outdoor sama seketi urut
11	negatif	kepal caran dora dora belum bedil outdoor kurang besar jd warna kurang mampu paku beli anda outdoor jd mungkin warna kurang mampu beli bener bel
12	negatif	mungkin outdoor beli anah enagang petakah laka kama guna enagang seketi hapus negasi n mungkin outdoor enagang seketi minimal rasa marah jengal/led saudara bangsa
13	positif	kampung nengang baret bangat bodi kama cewe paku motor
14	negatif	banan kenderan budil bangrak rekam batar rekam amekia rasa masah bader kenderan batar
15	positif	harga telaku mahal

Gambar 3. Hasil *Preprocessing*

4.1.3. Data Transformation

Gambar 4. Operator *Nominal to Text*

Proses selanjutnya adalah transformasi data. Operator *Nominal to text* untuk menyesuaikan jenis data yang akan digunakan yaitu atribut text.

Gambar 5. Operator *Process Documents from Data*

Operator *Process Documents from Data* membutuhkan atribut text dapat dilihat pada gambar 5 untuk memproses semua data menjadi satu dokumen. Terdapat beberapa operator didalamnya yaitu :

1. *Tokenize*
Operator ini digunakan untuk proses pemecahan teks menjadi bagian kecil yang disebut token sehingga dapat dianalisa.
2. *Stopword*
Penggunaan stopwords merupakan suatu langkah dalam mengeliminasi kata-kata yang tidak relevan dengan topik dokumen, bertujuan untuk meningkatkan *accuracy* dalam proses pengklasifikasian. *Stopword* ini berfungsi untuk menghapus kata-kata yang tidak memiliki makna signifikan ketika digunakan secara mandiri.
3. *Filter Token by Length*
Adalah proses menghapus kata tidak mengandung makna. Menggunakan parameter paling sedikit 4 huruf dan paling banyak 25 (*default*)
4. *Stem (Dictionary)*
Operator terakhir pada proses documents from data ini digunakan untuk memproses pemetaan kata dari bentuk kata menjadi bentuk kata dasar.

Proses selanjutnya adalah pembobotan kata. Pada proses ini menggunakan metode *TF-IDF* dimana setiap frekuensi kata akan dihitung kemunculannya pada data komentar *youtube*. Metode *Term Frequency-Inverse Document Frequency (TF-IDF)* adalah suatu teknik transformasi yang mengukur signifikansi kata dalam dokumen dengan memberikan nilai bobot pada masing-masing kata. Contoh kata “*handphone*” dengan nilai 0 artinya kata tersebut tidak muncul pada dokumen pertama sedangkan kata “*harga*” bernilai 0.094 berarti kata tersebut muncul pada dokumen pertama dan dokumen lainnya.

Gambar 6. Hasil Proses Pembobotan Kata TF-IDF

Berikut dokumen yang akan digunakan dalam perhitungan dengan metode TF-IDF.

Tabel 1. Dokumen yang digunakan dalam perhitungan TF-IDF

Dokumen	Teks
Dokumen 1 (D1)	Saran sih bikin ionic sama kaya brio insya alloh laris manis
Dokumen 2 (D2)	Problem kualitas diturunin harga dinaikin usaha gitu cari cuan subsidi sebab inflasi paling gede

Perhitungan TF dilakukan pada 2 sampel dokumen. Kata yang muncul diberi nilai 1 dan jika kata tidak muncul diberi nilai 0. Berikut hasil perhitungan TF (Term Frequency) dapat dilihat pada tabel berikut.

Tabel 2. Hasil Perhitungan TF

Tokenize	TF			
	D1	D2	D3	D4
saran	1	0	0	0
ionic	1	0	0	0
kaya	1	0	0	0
brio	1	0	0	0
insya	1	0	0	0
alloh	1	0	0	0
laris	1	0	0	0
manis	1	0	0	0
problem	0	1	0	0
subsidi	0	1	0	0
kualitas	0	1	2	1
turun	0	1	0	0
harga	0	1	0	1
naik	0	1	0	0
usaha	0	1	0	0
cuan	0	1	0	0
inflasi	0	1	0	0
gede	0	1	0	0

Kemudian dari hasil perhitungan TF (Term Frequency). Selanjutnya menghitung jumlah dokumen IDF, IDF ialah hasil inverse dari DF

Tabel 3. Hasil Perhitungan IDF

DF	D/Df	IDF (LogD/Df)
1	4	0,6
1	4	0,6
1	4	0,6
1	4	0,6
1	4	0,6

DF	D/Df	IDF (LogD/Df)
1	4	0,6
1	4	0,6
1	4	0,6
1	4	0,6
1	4	0,6
3	1,33	0,6
1	4	0,6
2	2	0,6
1	4	0,6
1	4	0,6
1	4	0,6
1	4	0,6
1	4	0,6

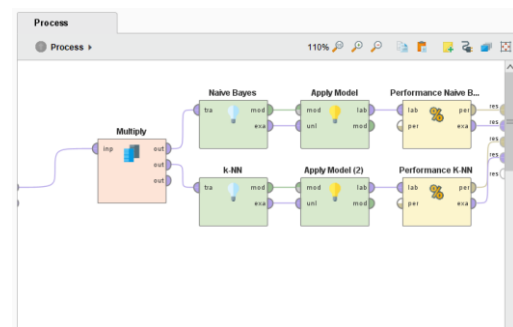
Berikut dapat dihitung hasil TF-IDF dengan mengalikan hasil TF dan IDF. Hasilnya dapat dilihat pada tabel dibawah ini.

Tabel 4. Hasil Perhitungan TF-IDF

W			
D1	D2	D3	D4
0,6	0	0	0
0,6	0	0	0
0,6	0	0	0
0,6	0	0	0
0,6	0	0	0
0,6	0	0	0
0,6	0	0	0
0,6	0	0	0
0	0,6	0	0
0	0,6	0	0
0	0,12	0,24	0,12
0	0	0	0,6
0	0,3	0	0,3
0	0,6	0	0
0	0,6	0	0
0	0,6	0	0
0	0,6	0	0
0	0,6	0	0

4.1.4. Data Mining

Data mining adalah proses mengekstraksi pola atau pengetahuan yang berharga dari sekumpulan data besar. Penelitian ini membagi data training sebanyak 0.8 dan data testing sebanyak 0.2 pada parameter operator *Split Data*. Kemudian operator *Multiply* digunakan untuk membuat salinan objek.



Gambar 7. Proses Data Mining

1. Naïve Bayes

accuracy: 86,95%				
	true positif	true negatif	true netral	class precision
pred positif	368	80	0	82,14%
pred negatif	0	571	0	100,00%
pred netral	34	44	114	59,38%
class recall	91,54%	82,10%	100,00%	

Gambar 8. Hasil Naïve Bayes

Berdasarkan hasil Performance Naïve Bayes mengenai analisis sentimen pada komentar tentang intensif mobil listrik di laman youtube diperoleh tingkat *accuracy* sebesar 86,95%, dengan nilai *precision* untuk kategori positif sebesar 82,14%, untuk kategori negatif 100,00% dan untuk kategori netral sebesar 59,38%. Lalu nilai *recall* untuk kategori positif sebesar 91,54%, untuk kategori negatif 82,16% dan kategori netral sebesar 100,00%.

2. K-Nearest Neighbor

accuracy: 92,23%				
	true positif	true negatif	true netral	class precision
pred positif	400	49	9	87,34%
pred negatif	2	644	20	96,70%
pred netral	0	2	85	87,70%
class recall	99,50%	92,66%	74,56%	

Gambar 9. Hasil K-Nearest Neighbor

Berdasarkan hasil *performance* mengenai analisis sentimen pada komentar tentang intensif mobil listrik di laman youtube menggunakan metode K-Nearest Neighbor diperoleh hasil tingkat *accuracy* sebesar 92,23%, dengan nilai *precision* untuk kategori positif sebesar 87,34%, untuk kategori negatif 96,70% dan untuk kategori netral sebesar 87,70%. Lalu nilai *recall* untuk kategori positif sebesar 99,50%, untuk kategori negatif 92,66% dan kategori netral sebesar 74,56%.

4.1.5. Patern Evaluation

Setelah proses klasifikasi dengan Naïve Bayes dan K-Nearest Neighbor. Tahap berikutnya ialah melakukan proses evaluasi yang bertujuan untuk mengetahui nilai *accuracy*, *precision*, dan *recall* dari hasil yang didapatkan. *Confusion Matrix* digunakan dalam evaluasi ini karena berperan sebagai sarana evaluasi untuk mengestimasi objek yang teridentifikasi dengan benar (*True*) dan salah (*False*) pada model klasifikasi. Berikut adalah *Confusion Matrix* dapat dilihat pada tabel 6 di bawah ini:

Tabel 5. Confusion Matrix Tiga Kelas

		Kelas Aktual		
		Positif	Negatif	Netral
Kelas Prediksi	Positif	True Positive (TP)	False Negative (FN)	False Neutral (FN)
	Negatif	False Positive (FP)	True Negative (TN)	False Neutral (FN)
	Netral	False Neutral (FN)	False Neutral (FN)	True Netral (TN)

Rumus dan perhitungan untuk mencari nilai *accuracy* Naïve Bayes :

$$\begin{aligned} \text{Accuracy} &= \frac{TP_{\text{Positif}} + TN_{\text{Negatif}} + TN_{\text{Netral}}}{TP_{\text{Positif}} + (FN_{\text{Negatif}} + FN_{\text{Netral}})} \times 100\% \\ &= \frac{368 + 571 + 114}{368 + 571 + 114} \times 100\% \\ &= \frac{1.053}{1.211} \times 100\% = 86,95\% \end{aligned} \quad (1)$$

Rumus dan perhitungan untuk mencari nilai *precision* positif Naïve Bayes

$$\begin{aligned} \text{Precision Positif} &= \frac{TP_{\text{Positif}}}{TP_{\text{Positif}} + (FN_{\text{Negatif}} + FN_{\text{Netral}})} \times 100\% \\ \text{Precision Positif} &= \frac{368}{368 + 80} \times 100\% \\ &= \frac{368}{448} \times 100\% = 82,14\% \end{aligned} \quad (2)$$

Rumus dan perhitungan untuk mencari nilai *precision* negatif Naïve Bayes

$$\begin{aligned} \text{Precision Negatif} &= \frac{TN_{\text{Negatif}}}{TN_{\text{Negatif}} + (FP_{\text{Positif}} + FN_{\text{Netral}})} \times 100\% \\ \text{Precision Negatif} &= \frac{571}{571 + 0} \times 100\% \\ &= \frac{571}{571} \times 100\% = 100,00\% \end{aligned} \quad (3)$$

Rumus dan perhitungan untuk mencari nilai *precision* netral Naïve Bayes

$$\begin{aligned} \text{Precision Netral} &= \frac{TN_{\text{Netral}}}{TN_{\text{Netral}} + (FN_{\text{Netral}} + FN_{\text{Netral}})} \times 100\% \\ \text{Precision Netral} &= \frac{114}{114 + 78} \times 100\% \\ &= \frac{114}{192} \times 100\% = 59,38\% \end{aligned} \quad (4)$$

Rumus dan perhitungan untuk mencari nilai *recall* positif Naïve Bayes

$$\begin{aligned} \text{Recall Positif} &= \frac{TP_{\text{Positif}}}{TP_{\text{Positif}} + (FP_{\text{Positif}} + FN_{\text{Netral}})} \times 100\% \\ \text{Recall Positif} &= \frac{368}{368 + 34} \times 100\% \\ &= \frac{368}{402} \times 100\% = 91,54\% \end{aligned} \quad (5)$$

Rumus dan perhitungan untuk mencari nilai *recall* negatif Naïve Bayes

$$\begin{aligned} \text{Recall Negatif} &= \frac{TN_{\text{Negatif}}}{TN_{\text{Negatif}} + (TN_{\text{Negatif}} + FN_{\text{Netral}})} \times 100\% \\ \text{Recall Negatif} &= \frac{571}{571 + 124} \times 100\% \\ &= \frac{571}{695} \times 100\% = 82,16\% \end{aligned} \quad (6)$$

Rumus dan perhitungan untuk mencari nilai *recall* netral Naïve Bayes

$$\begin{aligned} \text{Recall Netral} &= \frac{TN_{\text{Netral}}}{TN_{\text{Netral}} + (TN_{\text{Netral}} + FN_{\text{Netral}})} \times 100\% \\ \text{Recall Netral} &= \frac{114}{114 + 0} \times 100\% \\ &= \frac{114}{114} \times 100\% = 100,00\% \end{aligned} \quad (7)$$

Rumus dan perhitungan untuk mencari nilai *accuracy* K-Nearest Neighbor:

$$\begin{aligned} \text{Accuracy} &= \frac{TP_{\text{Positif}} + TN_{\text{Negatif}} + TN_{\text{Netral}}}{\text{Jumlah Data Keseluruhan}} \times 100\% \\ \text{Accuracy} &= \frac{400 + 644 + 85}{1.229} \times 100\% \\ &= \frac{1.129}{1.211} \times 100\% = 93,23\% \end{aligned} \quad (8)$$

Rumus dan perhitungan untuk mencari nilai *precision* positif K-Nearest Neighbor

$$\begin{aligned} \text{Precision Positif} &= \frac{TP_{\text{Positif}}}{TP_{\text{Positif}} + (FN_{\text{Negatif}} + FN_{\text{Netral}})} \times 100\% \\ \text{Precision Positif} &= \frac{400}{400 + 58} \times 100\% \\ &= \frac{400}{458} \times 100\% = 87,34\% \end{aligned} \quad (9)$$

Rumus dan perhitungan untuk mencari nilai *precision* negatif K-Nearest Neighbor

$$\begin{aligned} \text{Precision Negatif} &= \frac{TN_{\text{Negatif}}}{TN_{\text{Negatif}} + (FP_{\text{Positif}} + FN_{\text{Netral}})} \times 100\% \\ \text{Precision Negatif} &= \frac{644}{644 + 22} \times 100\% \\ &= \frac{644}{666} \times 100\% = 96,70\% \end{aligned} \quad (10)$$

Rumus dan perhitungan untuk mencari nilai *precision* netral K-Nearest Neighbor

$$\begin{aligned} \text{Precision Netral} &= \frac{TN_{\text{Netral}}}{TN_{\text{Netral}} + (FN_{\text{Netral}} + FN_{\text{Netral}})} \times 100\% \\ \text{Precision Netral} &= \frac{85}{85 + 2} \times 100\% \\ &= \frac{85}{87} \times 100\% = 97,70\% \end{aligned} \quad (11)$$

Rumus dan perhitungan untuk mencari nilai *recall* positif *K-Nearest Neighbor*

$$\begin{aligned} \text{Recall Positif} &= \frac{TP_{\text{Positif}}}{\frac{400}{TP_{\text{Positif}} + (FP_{\text{Netral}})}} \times 100\% \\ \text{Recall Positif} &= \frac{400}{400 + 2} \times 100\% \\ &= \frac{400}{402} \times 100\% = \mathbf{99,50\%} \end{aligned} \quad (12)$$

Rumus dan perhitungan untuk mencari nilai *recall* negatif *K-Nearest Neighbor*

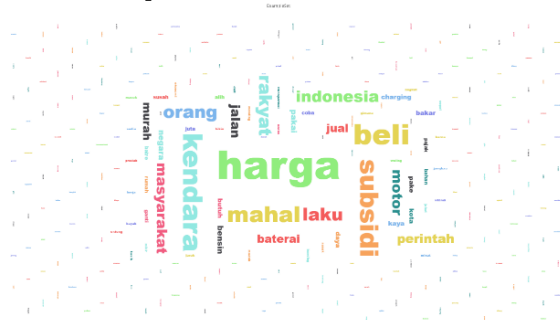
$$\begin{aligned} \text{Recall Negatif} &= \frac{TNegatif}{TNegatif + (TNegatif + FNutral)} \times 100\% \\ \text{Recall Negatif} &= \frac{644}{644 + 51} \times 100\% \\ &= \frac{644}{695} \times 100\% = \mathbf{92,66\%} \end{aligned} \quad (13)$$

Rumus dan perhitungan untuk mencari nilai *recall* netral *K-Nearest Neighbor*

$$\begin{aligned} \text{Recall Netral} &= \frac{TNetral}{TNetral + (TNetral + FNetral)} \times 100\% \\ \text{Recall Netral} &= \frac{85}{85 + 29} \times 100\% \\ &= \frac{85}{114} \times 100\% = \mathbf{74,56\%} \end{aligned} \quad (14)$$

4.1.6. Knowledge Presentation

Hasil preprocessing data dikelompokkan dengan tujuan mengidentifikasi frekuensi kemunculan kata-kata yang tinggi. Informasi frekuensi kata-kata ini akan direpresentasikan secara visual melalui pembuatan *wordcloud*.



Gambar 10. Hasil Visualisasi *Wordcloud*

Dalam gambar 10 frekuensi kemunculan kata yang paling tinggi adalah kata “harga” dengan total kemunculan sebesar 446 kata, disusul oleh kata “beli” dengan total kemunculan sebesar 325 kata. Kata yang muncul terlihat lebih kecil maka kata tersebut berarti nilai frekuensi kemunculan katanya lebih rendah.

4.2. Pembahasan Penelitian

Penerapan algoritma Naïve Bayes dan K-Nearest Neighbor (K-NN). Penelitian ini menerapkan dua algoritma *Naïve Bayes* dan *K-Nearest Neighbor*. Peneliti mengamati, berdasarkan jurnal penelitian-penelitian terdahulu yang pernah dibaca oleh peneliti metode *Naïve Bayes* sangat populer digunakan karena performa akurasi yang sangat baik dan tepat dalam analisis sentimen. Algoritma *Naïve Bayes* dipilih karena keunggulannya dalam proses klasifikasi, dimana hanya membutuhkan sedikit data dari data pelatihan untuk memilih parameter yang diperlukan, juga prosesnya cepat dan efisien [9].

K-Nearest Neighbor juga dipilih karena kelebihanannya dalam menangani data pelatihan dengan jumlah noise yang tinggi, cepat, mudah dipahami dan

kemampuan untuk menangani data dalam jumlah besar [10]. Algoritma ini terbukti efektif, seperti yang dibuktikan dalam penelitian “Implementasi Algoritma *Naive Bayes*, *Support Vector Machine*, dan *K-Nearest Neighbors* Untuk Analisa Sentimen Aplikasi Halodoc” (Indrayuni dkk, 2021)

Hasil penerapan algoritma *Naïve Bayes* dan *K-Nearest Neighbor* didapat hasil klasifikasi adalah sebagai berikut :

Tabel 6. Hasil Klasifikasi Ulasan Sentimen *Youtube*

Algoritma	Prediksi Sentimen		
	Positif	Negatif	Netral
Naïve Bayes	368	571	114
K-Nearest Bayes	400	644	85

Setelah proses klasifikasi dengan *Naïve Bayes* dan *K-Nearest Neighbor*. Menghasilkan nilai confusion matrix berikut ini :

1. Naïve Bayes

Tabel 7. *Confusion Matrix Naïve Bayes*

		Kelas Aktual		
		Positif	Negatif	Netral
Kelas Prediksi	Positif	368	80	0
	Negatif	0	571	0
	Netral	34	44	144

Setelah dilakukan pengujian perhitungan manual menggunakan rumus confusion matrix tabel diatas berhasil mengklasifikasikan secara benar dengan nilai (True Positif) sebesar 368, data (True Negatif) 571, dan (True Netral) sebesar 144. Adapun hasil dari perhitungan nilai accuracy, precision dan recall terdapat pada tabel dibawah ini.

Tabel 8. Hasil Perhitungan Manual *Confusion Matrix* *Naïve Bayes*

	Precision	Recall	Accuracy
Positif	82,14%	91,54%	86,95%
Negatif	100,00%	82,16%	
Netral	59,38%	100,00%	

2. *K*-Nearest Neighbor

Tabel 9. *Confusion Matrix K-Nearest Neighbor*

		Kelas Aktual		
		Postif	Negatif	Netral
Kelas Prediksi	Positif	400	49	9
	Negatif	2	644	20
	Netral	0	2	85

Setelah dilakukan pengujian perhitungan manual menggunakan rumus *confusion matrix* tabel diatas berhasil mengklasifikasikan secara benar dengan nilai (*True Positif*) sebesar 400, data (*True Negatif*) 644, dan (*True Netral*) sebesar 85. Adapun hasil dari perhitungan nilai *accuracy*, *precision* dan *recall* terdapat pada tabel dibawah ini.

Tabel 10. Hasil Perhitungan Manual *Confusion Matrix K-Nearest Neighbor*

	Precision	Recall	Accuracy
Positif	87,34%	99,50%	93,23%
Negatif	96,70%	92,66%	
Netral	97,70%	74,56%	

Hasil Perbandingan Algoritma *Naïve Bayes* dan *K-Nearest Bayes*

Tabel 11. Hasil Perbandingan Algoritma *Naïve Bayes* dan *K-Nearest Neighbor*

Metode	Accuracy	Precision	Recall
<i>Naïve Bayes</i>	86,95%	80,51%	91,23%
<i>K-Nearest Neighbor</i>	93,23%	93,91%	91,56%

Hasil *accuracy* metode *Naïve Bayes* memperoleh hasil sebesar 86,95% dan metode *K-Nearest Neighbor* memperoleh hasil sebesar 93,23%. Sehingga penulis dapat menyimpulkan dari hasil *accuracy* perbandingan pada penelitian ini klasifikasi *K-Nearest Neighbor* merupakan metode klasifikasi terbaik dalam analisis sentimen mengenai intensif mobil listrik dengan hasil prediksi yang lebih akurat dengan tingkat *accuracy* sebesar 93,23%.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil analisis yang dilakukan terkait dengan sentimen intensif mobil listrik sebanyak 1517 data komentar pada laman *youtube* menggunakan algoritma *Naïve Bayes* dan *K-Nearest Neighbor* dan pengujian menggunakan *confussion matrix* algoritma *Naïve Bayes*, *accuracy* yang diperoleh sebesar 86,95% dan algoritma *K-Nearest Neighbor* memperoleh hasil *accuracy* sebesar 93,23%. Maka dapat disimpulkan bahkan metode *K-Nearest Neighbor* lebih unggul dalam mengklasifikasi *text mining* dengan hasil perolehan *accuracy* 93,23, *precision* 93,91 dan *recall* 91,56. Saran untuk penelitian selanjutnya adalah diharapkan menggunakan dataset lebih banyak mengenai fenomena yang sedang tren. Menerapkan metode algoritma yang tidak banyak peneliti gunakan sehingga dapat mengetahui ferforma metode-metode yang digunakan dan menggunakan sumber data selain *youtube*.

DAFTAR PUSTAKA

- [1] S. Wegner, F. Mersmann, and M. G. Grados, "G20 Response To the Energy Crisis: Critical for 1.5°C," 2022. [Online]. Available: www.climate-transparency.org
- [2] S. Alfarizi and E. Fitriani, "Analisis Sentimen Kendaraan Listrik Menggunakan Algoritma Naive Bayes dengan Seleksi Fitur Information Gain dan Particle Swarm Optimization," *Indones. J. Softw. Eng.*, vol. 9, no. 1, pp. 19–27, 2023, [Online]. Available: <http://ejournal.bsi.ac.id/ejurnal/index.php/ijse>
- [3] A. F. Riyadi, F. R. Rahman, M. A. Nofa Pratama, M. K. Khafidli, and H. Patria, "Pengukuran Sentimen Sosial Terhadap Teknologi Kendaraan Listrik: Bukti Empiris di Indonesia," *Expert J. Manaj. Sist. Inf. dan Teknol.*, vol. 11, no. 2, p. 141, 2021, doi: 10.36448/expert.v11i2.2171.
- [4] E. Indrayuni, A. Nurhadi, and D. A. Kristiyanti, "Implementasi Algoritma Naive Bayes, Support Vector Machine, dan K-Nearest Neighbors untuk Analisa Sentimen Aplikasi Halodoc," *Fakt. Exacta*, vol. 14, no. 2, p. 64, 2021, doi: 10.30998/faktorexacta.v14i2.9697.
- [5] A. Firdaus and W. I. Firdaus, "Text Mining Dan Pola Algoritma Dalam Penyelesaian Masalah Informasi: (Sebuah Ulasan)," *J. JUPITER*, vol. 13, no. 1, p. 66, 2021.
- [6] Pristiyono, M. Ritonga, M. A. Al Ihsan, A. Anjar, and F. H. Rambe, "Sentiment analysis of Covid-19 vaccine in Indonesia using Naïve Bayes Algorithm," *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 1088, no. 1, p. 012045, 2021, doi: 10.1088/1757-899x/1088/1/012045.
- [7] L. Ardiantoro, S. Zahara, and N. Sunarmi, "Pemanfaatan Knowledge Data Discovery(KDD) Pada Pola Permainan Atlet Bulutangkis," *Explor. IT J. Keilmuan dan Apl. Tek. Inform.*, vol. 11, no. 1, pp. 1–6, 2019, doi: 10.35891/explorit.v11i1.1467.
- [8] S. Alam, M. G. Resmi, and N. Masripah, "Classification of Covid-19 vaccine data screening with Naive Bayes algorithm using Knowledge Discovery in database method," *J. Comput. Networks, Archit. High Perform. Comput.*, vol. 4, no. 2, pp. 177–185, 2022, doi: 10.47709/cnahpc.v4i2.1584.
- [9] A. Erfina and M. R. N. R. Alamsyah, "Implementation of Naive Bayes classification algorithm for Twitter user sentiment analysis on ChatGPT using Python programming language," *Data Metadata*, vol. 2, pp. 2–11, 2023, doi: 10.56294/dm202345.
- [10] I. Tri Julianto, D. Kurniadi, Y. Septiana, and A. Sutedi, "Alternative Text Pre-Processing using Chat GPT Open AI," *J. Nas. Pendidik. Tek. Inform.*, vol. 12, no. 1, pp. 67–77, 2023, doi: 10.23887/janapati.v12i1.59746.