



Klasifikasi Sentimen Terhadap Kebijakan PHK 55 Ribu Karyawan oleh BT Group menggunakan *Algoritma Klasifikasi Naive Bayes*

**Muhammad Obie Charisma¹, Muhammad Farid Hamzah², M. Erwin³,
Intan Nurbaiti⁴, Fandi Kurniawan⁵**

^{1,2,3,4,5}Jurusan Sistem dan Teknologi Informasi, Universitas Muhammadiyah Kotabumi,
Lampung Utara, Indonesia

Email: muham.2059201107@umko.ac.id¹, muham.2059201112@umko.ac.id²,
merwi.2059201111@umko.ac.id³, intan.2059201080@umko.ac.id⁴,
fandi.kurniawan@umko.ac.id⁵

Abstrak

Artikel penelitian ini membahas analisis sentimen terkait kebijakan PHK yang dilakukan oleh BT Group terhadap 55 ribu orang karyawan sebagai dampak dari penggunaan kecerdasan buatan (AI). Metode Naive Bayes digunakan untuk melakukan klasifikasi sentimen pada komentar-komentar yang muncul pada platform YouTube terkait kebijakan tersebut. Tujuan utama dari analisis ini adalah untuk memahami respons dan sentimen yang muncul dari masyarakat terhadap kebijakan PHK yang diambil oleh BT Group. Dengan menggunakan metode Naive Bayes, artikel ini bertujuan untuk memberikan wawasan yang mendalam tentang bagaimana teknologi kecerdasan buatan dapat memengaruhi kebijakan perusahaan terkait PHK, serta respons masyarakat terhadap kebijakan tersebut berdasarkan analisis sentimen yang dilakukan. Dengan demikian, artikel ini diharapkan dapat memberikan kontribusi dalam pemahaman terhadap dampak penggunaan teknologi AI dalam konteks kebijakan perusahaan, serta respons masyarakat terhadap kebijakan tersebut berdasarkan analisis sentimen yang dilakukan.

Kata Kunci: PHK, Analisis Sentimen, Naïve Bayes

1. PENDAHULUAN

Pertumbuhan teknologi dan kecerdasan buatan (AI) telah mengubah paradigma dalam berbagai sektor, termasuk dunia bisnis. Perusahaan-perusahaan besar mulai mengadopsi kebijakan berbasis AI untuk meningkatkan efisiensi operasional dan produktivitas. Salah satu contohnya adalah BT Group, yang memutuskan untuk



mengimplementasikan teknologi AI dalam pengelolaan sumber daya manusianya. Namun, penggunaan AI tidak selalu berjalan tanpa kontroversi. Sebagai contoh, kebijakan PHK (Pemutusan Hubungan Kerja) massal yang dilakukan oleh BT Group, menciptakan getaran besar di masyarakat. Keputusan perusahaan untuk mem-PHK 55 ribu karyawan dipicu oleh penggunaan AI dalam pengambilan keputusan terkait sumber daya manusia.

Dalam artikel ini, kami akan melakukan analisis sentimen terhadap kebijakan BT Group menggunakan metode Naive Bayes. Analisis sentimen adalah riset komputasional dari opini, sentimen dan emosi yang diekspresikan secara tekstual [1]. Fokus analisis sentimen kami tidak hanya terbatas pada laporan dan berita konvensional, melainkan juga merambah ke dunia maya, khususnya komentar-komentar yang muncul di platform video seperti YouTube. Mengapa kita memilih YouTube? Karena YouTube telah menjadi salah satu platform utama di mana masyarakat berbagi pandangan mereka. Komentar-komentar yang ditinggalkan oleh pengguna memberikan gambaran langsung tentang bagaimana kebijakan PHK BT Group direspon oleh publik secara luas .

Metode yang kami gunakan dalam artikel analisis sentimen ini adalah metode naive bayes. Metode Naive Bayes adalah metode klasifikasi yang berdasarkan pada teorema Bayes. Metode ini digunakan untuk mengklasifikasikan data ke dalam kategori tertentu berdasarkan kemungkinan terjadinya suatu kejadian [2].

Penggunaan algoritma Klasifikasi Naive Bayes dalam konteks ini menjadi sangat penting, karena dapat membantu mengidentifikasi polaritas sentimen secara otomatis, yaitu apakah sentimen bersifat positif, negatif, atau netral. Tak hanya sebatas itu, algoritma ini juga memungkinkan kita untuk mengeksplorasi nuansa dan argumen di balik setiap sentimen yang muncul. Dengan demikian, penelitian ini diharapkan dapat memberikan kontribusi yang berarti dalam pemahaman dampak kebijakan PHK terhadap citra perusahaan dan persepsi masyarakat. Dengan pengetahuan ini, perusahaan dapat lebih baik memahami respons masyarakat dan mengambil langkah-langkah yang lebih berkelanjutan dalam pengambilan keputusan strategis mereka [3].

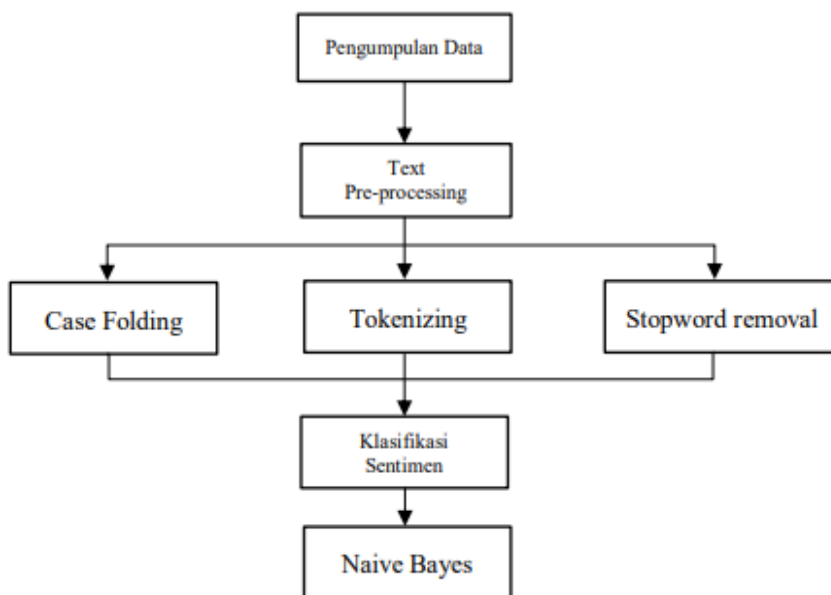
Metode Naive Bayes sangat efektif dalam melakukan analisis sentimen karena dapat mengklasifikasikan data ke dalam kategori positif, negatif,

atau netral [4]. Dengan menggunakan metode Naive Bayes, kita dapat mengklasifikasikan sentimen-sentimen yang muncul dalam komentar-komentar tersebut. Dari sini, kita dapat mengidentifikasi pola-pola umum, tren, dan tanggapan dominan yang muncul di kalangan pemirsa YouTube terhadap kebijakan kontroversial ini [5].

Analisis ini tidak hanya akan memberikan wawasan mendalam tentang bagaimana masyarakat merespons kebijakan PHK BT Group, tetapi juga memberikan pandangan tentang bagaimana teknologi AI dapat menjadi faktor yang mempengaruhi kebijakan sumber daya manusia dan dampaknya pada ketenangan sosial.

2. METODE

Dalam penelitian ini objek yang diteliti oleh penulis adalah kebijakan PHK 55 Ribu Karyawan oleh BT Group di media sosial Youtube. Data yang digunakan berupa komentar pada video yang berjudul kebijakan PHK 55 Ribu Karyawan oleh BT Group. Dalam penelitian ini, Tool Rapidminer digunakan untuk menjalankan serangkaian langkah metodologi guna mengumpulkan dan menganalisis data [6].



Gambar 1. Alur Penelitian

1) Pengumpulan Data

Pengumpulan Data: Proses pengumpulan data dilaksanakan menggunakan python dengan memanfaatkan netlytic. Peneliti melakukan ekstraksi terhadap 1.548 data komentar yang terdapat pada video YouTube berjudul "BT Group PHK 55 ribu orang karena pakai AI." Seluruh data yang berhasil diekstraksi akan diolah pada tahap selanjutnya melalui penggunaan tools Rapidminer.

2) Text Preprocessing

Text Preprocessing: Tahap preprocessing merupakan langkah kritis di mana dilakukan pemilihan data untuk membuatnya lebih terstruktur. Pada penelitian ini, tahap text preprocessing melibatkan beberapa proses, termasuk: case folding, tokenizing, dan filtering [7].

- a) Letter casing (case folding): Proses ini melibatkan konversi seluruh teks dalam komentar ke bentuk standar, yaitu huruf kecil atau lowercase. Hal ini dilakukan untuk memastikan konsistensi dalam analisis, di mana perbedaan huruf besar dan kecil diabaikan.
- b) Tokenizing: Pada tahap ini, komentar diubah menjadi token, yaitu kata-kata yang dipisahkan oleh spasi dalam teks. Proses ini bertujuan untuk membongkar teks menjadi unit-unit kata yang dapat diolah lebih lanjut.
- c) Stopword removal: Proses ini melibatkan penghapusan kata-kata yang dianggap tidak memberikan kontribusi signifikan pada analisis, seperti kata penghubung "dan," "atau," "kemudian," dan sejenisnya. Penghapusan stopwords bertujuan untuk memfokuskan perhatian pada kata-kata yang lebih berarti dan berpotensi memengaruhi proses klasifikasi.

Melalui tahap-tahap tersebut, data yang awalnya tidak terstruktur akan mengalami penyaringan dan pemrosesan sehingga lebih siap untuk tahapan analisis sentimen menggunakan algoritma klasifikasi Naive Bayes.

3) Klasifikasi

Klasifikasi: Setelah melalui proses preprocessing, komentar akan menjalani tahap klasifikasi sesuai dengan kelasnya (sentiment class) untuk menentukan polaritas teks tersebut, apakah termasuk dalam opini positif, negatif, atau netral. Proses klasifikasi ini menggunakan algoritma Klasifikasi Naive Bayes untuk mengidentifikasi dan mengategorikan sentimen yang terkandung dalam setiap komentar. Dengan demikian, tujuan utama dari tahap klasifikasi ini adalah untuk memberikan

pemahaman yang lebih mendalam terhadap sikap dan pandangan masyarakat terkait kebijakan PHK BT Group yang telah diungkapkan melalui komentar-komentar tersebut..

4) Penerapan metode Naive Bayes

Data yang telah diklasifikasikan sentimennya akan melewati tahap cross-validation, suatu proses yang bertujuan untuk mendapatkan nilai akurasi menggunakan metode Naive Bayes. Pada tahap cross-validation, dataset dibagi menjadi subset-subset kecil yang disebut lipatan (folds). Algoritma Naive Bayes kemudian diterapkan pada beberapa kombinasi subset training dan testing secara bergantian.

3. HASIL DAN PEMBAHASAN

3.1 *Crawling Data*

Crawling data, Mengumpulkan dataset yang relevan dengan studi kasus, seperti video youtube yang berjudul kebijakan bt group mem phk 55 ribu orang karena pakai ai menggunakan netlytic.



Gambar 2. Video Youtube yang akan diproses

3.2 Pre-Processing

Pre-processing data merupakan fase kritis dalam siklus analisis data, di mana langkah-langkah penting diambil untuk memastikan bahwa data yang telah terkumpul bersih, konsisten, dan siap untuk dianalisis. Proses ini melibatkan beberapa tahapan, yang pertama adalah pembersihan data, di mana data yang tidak lengkap, duplikat, atau memiliki nilai yang tidak valid diidentifikasi dan dihapus untuk memastikan keakuratan analisis yang akan datang. Setelah itu, transformasi data dilakukan untuk mengubah format atau struktur data agar lebih sesuai dengan kebutuhan analisis. Encoding variabel kategorikal merupakan langkah selanjutnya, di mana variabel kategori diubah menjadi bentuk yang dapat diinterpretasikan oleh model atau algoritma. Penghapusan *noise* dilakukan untuk menghilangkan gangguan atau variabilitas yang tidak relevan dari data. Pengaturan fitur dan ekstraksi fitur kemudian dilibatkan untuk memastikan bahwa data yang digunakan memiliki representasi yang optimal dalam konteks analisis tertentu. Selama proses pembersihan data, *Google Colab* digunakan sebagai perangkat lunak yang andal untuk memastikan integritas dan kebersihan data, membantu membangun dasar yang kuat untuk analisis data yang efektif dan akurat [2].

	A
1	text
2	find batu mulia
3	jasa perbaikan mesin ai akan menjadi lahan pekerjaan baru
4	buat novel
5	sweater puisi
6	manusia akan digantikan dengan mesin
7	saham nvidia n oracle meroket
8	klo smua pke ai boker pun bakal dicebokin robot
9	mimpi buruk akan segera datang ke indonesia bersama dengan ai

Gambar 3. Data dalam bentuk Text

3.3 TF-IDF

Metode *Term Frequency Inverse Document Frequency (TF-IDF)* adalah suatu pendekatan yang digunakan untuk menilai seberapa penting suatu kata (term) dalam suatu dokumen dengan memberikan bobot pada masing-masing kata tersebut. Proses perhitungan bobot menggunakan *TF-IDF* melibatkan dua aspek utama, yaitu *Term Frequency (TF)* dan *Inverse Document Frequency (IDF)* [8]. Pertama-tama, nilai *TF* per kata

dihitung untuk menunjukkan seberapa sering kata tersebut muncul dalam suatu dokumen. Rumusnya dinyatakan sebagai berikut:

$$TF_{kata} = \frac{\text{Jumlah kemunculan kata dalam dokumen}}{\text{Jumlah total kata dalam dokumen}}$$

Selanjutnya, bobot masing-masing kata dihitung dengan menggunakan nilai *IDF*. *IDF* memberikan bobot pada suatu kata berdasarkan seberapa umum atau langka kata tersebut muncul di seluruh koleksi dokumen. Persamaan *IDF* dirumuskan sebagai berikut:

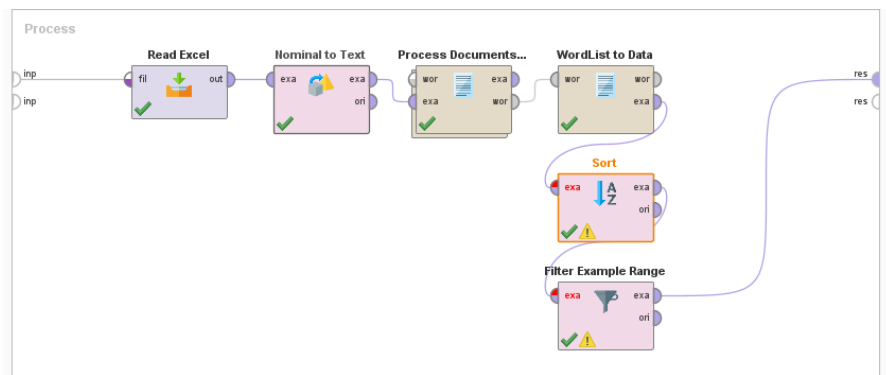
$$IDF_{kata} = \log \left(\frac{\text{Jumlah total dokumen dalam koleksi}}{\text{Jumlah dokumen yang mengandung kata}} \right)$$

Dalam rumus ini, *DF* (*\text{kata}*) mengacu pada jumlah kemunculan kata dalam seluruh dokumen dalam koleksi, sedangkan *td* adalah jumlah dokumen dalam koleksi. Gabungan dari nilai *TF* dan *IDF* menghasilkan bobot akhir suatu kata dalam suatu dokumen, yang dinyatakan sebagai:

$$TF-IDF_{kata} = TF_{kata} \times IDF_{kata}$$

Penerapan metode *TF-IDF* ini memiliki tujuan untuk menyoroti kata-kata yang memiliki relevansi tinggi terhadap suatu dokumen dalam konteks koleksi dokumen yang lebih besar. Dengan memberikan bobot pada kata-kata tersebut, *TF-IDF* membantu mengidentifikasi dan menekankan kata-kata kunci yang dapat memberikan kontribusi besar terhadap pemahaman dan representasi suatu dokumen dalam analisis teks.

Dengan menggunakan *DF* dan *IDF*, metode *TF-IDF* memberikan bobot pada kata-kata berdasarkan seberapa umum atau langka kata tersebut muncul di seluruh koleksi dokumen. Semakin sering kata tersebut muncul di berbagai dokumen, nilai *DF* (*\text{word}*) akan meningkat, tetapi bobot *IDF* akan menurun, dan sebaliknya. Kombinasi *TF* dan *IDF* kemudian menghasilkan nilai bobot *TF-IDF* yang mencerminkan tingkat pentingnya suatu kata dalam suatu dokumen terhadap seluruh koleksi dokumen [4].



Gambar 4. Tahapan TF-IDF

3.4 WordCloud

World Cloud tidak hanya sekadar menjadi kumpulan kata-kata dari sumber data, tetapi juga menjalani proses pengolahan dan transformasi yang mendalam untuk menghasilkan representasi visual yang informatif. Sebagai bagian integral dari *Natural Language Processing (NLP)*, teknik ini memanfaatkan Bahasa atau text mining untuk mendalami dan menganalisis permasalahan yang terkait dengan data teks. Dalam konteks ini, word cloud menjadi alat visualisasi yang sangat berguna. Setelah kata-kata diekstraksi, bobot atau frekuensi kemunculan masing-masing kata dianalisis, dan kata-kata yang dominan kemudian divisualisasikan dalam bentuk word cloud. Word cloud memberikan pandangan yang langsung dan tangkas mengenai kata-kata yang paling berpengaruh atau sering muncul dalam dataset, memungkinkan para analis untuk dengan cepat mengidentifikasi pola, fokus, atau tren utama yang terkandung dalam teks tersebut [9].

3.5 Visualisasi Word Cloud

Pada langkah ini, representasi hasil preprocessing data ditampilkan melalui word cloud visualisasi untuk setiap kelas sentimen. Word cloud ini menampilkan kata-kata dengan frekuensi tertinggi, di mana ukuran kata-kata mencerminkan seberapa sering kata-kata tersebut muncul dalam dataset. Semakin tinggi frekuensi kemunculan kata, semakin besar ukuran kata tersebut dalam word cloud, mengindikasikan bahwa kata tersebut lebih sering muncul dalam dataset yang digunakan. Visualisasi ini memberikan gambaran intuitif tentang kata-kata yang... (lanjutannya tergantung pada konteks selanjutnya) paling umum atau dominan dalam

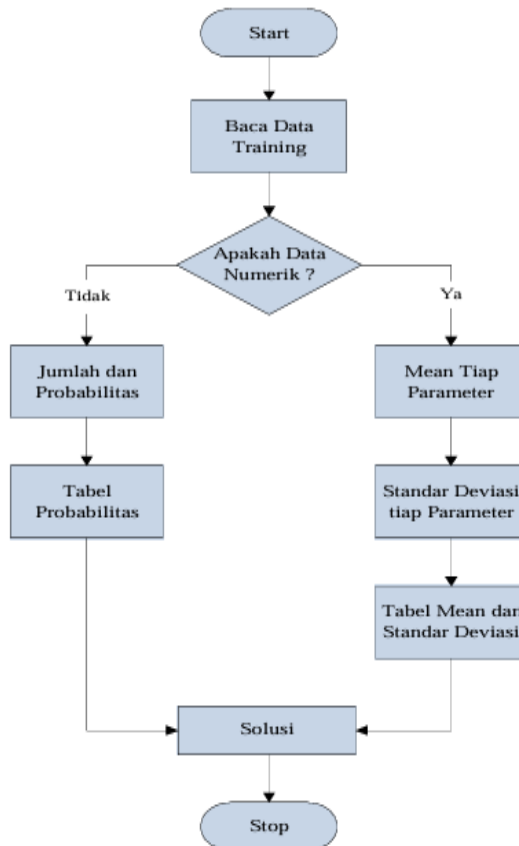
setiap kelas sentimen, membantu pemahaman terhadap karakteristik dan pola sentimen dalam data tersebut [10].



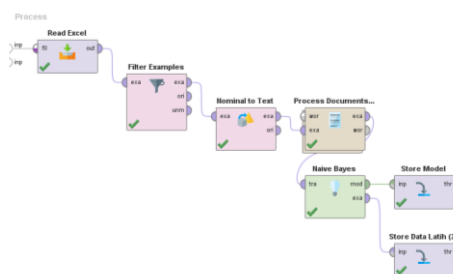
Gambar 5. Word Cloud

3.6 Naive Bayes Classifier

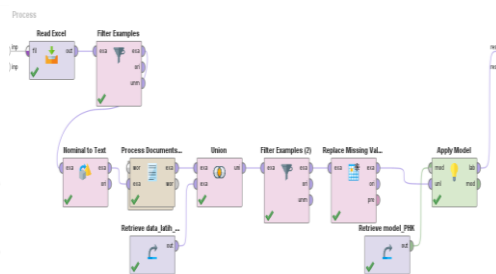
Naive Bayes Classifier, dalam esensinya, menyajikan suatu pendekatan yang relatif sederhana namun efektif dalam konteks pengklasifikasi probabilistik. Proses operasionalnya melibatkan perhitungan probabilitas yang berkaitan dengan sekelompok atribut dengan cara menjumlahkan frekuensi dan kombinasi nilai yang terkandung dalam dataset yang sedang dianalisis. Kunci keberhasilan algoritma ini terletak pada penerapan Teorema Bayes, yang memungkinkan penyesuaian probabilitas kelas berdasarkan informasi atribut yang diberikan. Penting untuk dicatat bahwa *Naive Bayes Classifier* mengasumsikan bahwa semua atribut yang diperhitungkan bersifat independen atau tidak saling ketergantungan, sebuah asumsi yang menjadi dasar bagi kemudahan implementasi dan kecepatan perhitungan. Hal ini sangat relevan dalam konteks analisis sentiment data training, di mana algoritma ini mampu memberikan hasil yang cukup akurat dengan mempertimbangkan nilai variabel kelas untuk mengklasifikasikan data dengan tepat [11].



Gambar 6. Skema Naive Bayes [5]



Gambar 7. Naive Bayes Classifier Operator



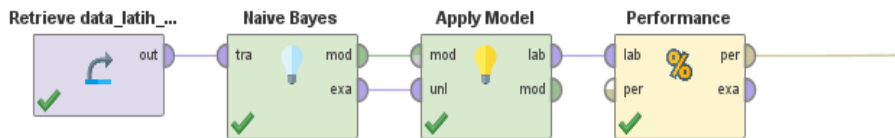
Gambar 8. Naive Bayes Classifier Operator

Row No.	Sentimen	prediction(S...	confidence(...	confidence(...	abadi	abal	abisan	accenture	accr
1	?	Positif	1	0	0	0	0	0	0
2	?	Positif	1	0	0	0	0	0	0
3	?	Positif	1	0	0	0	0	0	0
4	?	Positif	1	0	0	0	0	0	0
5	?	Positif	1	0	0	0	0	0	0
6	?	Negatif	0	1	0	0	0	0	0
7	?	Positif	1	0	0	0	0	0	0
8	?	Positif	1	0	0	0	0	0	0
9	?	Negatif	0	1	0	0	0	0	0
10	?	Positif	1	0	0	0	0	0	0
11	?	Negatif	0	1	0	0	0	0	0
12	?	Negatif	0	1	0	0	0	0	0
13	?	Positif	1	0	0	0	0	0	0
14	?	Positif	1	0	0	0	0	0	0

Gambar 9. Hasil dari Naive Bayes Classifier Operator

3.7 Performance

Kinerja (*performance*) dalam data mining adalah ukuran sejauh mana suatu sistem atau algoritma data mining mampu menjalankan tugasnya dengan efisien dan efektif. Kinerja ini dapat diukur dengan berbagai metrik yang sesuai dengan tujuan analisis data yang dilakukan. Beberapa faktor yang dapat memengaruhi kinerja dalam data mining meliputi akurasi hasil, kecepatan eksekusi, kemampuan untuk mengekstrak pola yang berarti, serta kemampuan untuk menangani data yang besar dan kompleks. Dalam mengevaluasi performa algoritma Naive Bayes Classifier, langkah selanjutnya melibatkan pencatatan nilai dan analisis menggunakan *Confusion Matrix*. *Confusion Matrix* memberikan gambaran yang terperinci tentang sejauh mana algoritma dapat mengklasifikasikan data dengan benar, dengan memperhitungkan *True Positives*, *True Negatives*, *False Positives*, dan *False Negatives*. Setelah mendapatkan *Confusion Matrix* dari berbagai percobaan, evaluasi dilakukan dengan menghitung nilai rata-rata dari berbagai metrik yang dapat diambil dari matriks tersebut, seperti akurasi, presisi, recall, dan *F1-score*. Proses ini memberikan pemahaman holistik tentang kinerja algoritma dalam berbagai konteks dan kondisi data. Nilai rata-rata yang dihasilkan dari evaluasi ini tidak hanya memberikan indikasi tentang seberapa baik model dapat memprediksi, tetapi juga membantu dalam mengidentifikasi percobaan yang dapat dijadikan acuan untuk penggunaan model algoritma yang dipilih. Dengan pendekatan ini, pengguna dapat membuat keputusan yang lebih informasional dan tepat terkait penerapan *Naive Bayes Classifier* pada data spesifik, meningkatkan kehandalan dan efektivitas model yang dikembangkan [3]. Berikut adalah *Confussion matrix*.



Gambar 9. Tahap Performance

accuracy: 93.67%

	true Positif	true Negatif	class precision
pred. Positif	420	0	100.00%
pred. Negatif	36	113	75.84%
class recall	92.11%	100.00%	

Gambar 10. Akurasi Naive Bayes Classifier

Metode naive Bayes berhasil memprediksi kebijakan bt group mem phk 55 ribu orang karena pakai ai dengan tingkat akurasi sebesar 93,67%. Dengan true positif sebesar 92,11%, true negatif sebesar 100,00%, class precision pred. Positif sebesar 100,00% dan class precision pred. negatif sebesar 75,84%.

4. KESIMPULAN

Berdasarkan penelitian tentang analisis sentimen kebijakan bt group mem phk 55 ribu orang karena pakai ai menggunakan metode *naive bayes* berdasarkan komentar pada youtube, maka dapat disimpulkan sebagai berikut: Metode *naive Bayes* dapat memakai data training untuk menciptakan probabilitas tiap kriteria dalam kelas yang berbeda. Dengan demikian, nilai probabilitas dari kriteria tersebut bisa dimaksimalkan dalam konteks kebijakan bt group mem phk 55 ribu orang karena pakai ai. Berdasarkan pemanfaatan komentar di video youtube yang berjudul kebijakan bt group mem phk 55 ribu orang karena pakai ai sebagai data training untuk metode *naive Bayes*, terdapat 1548 data yang dihasilkan, dan setelah proses ini, jumlahnya berkurang menjadi 1035. Oleh karena itu, metode naive Bayes berhasil memprediksi kebijakan bt group mem phk 55 ribu orang karena pakai ai dengan tingkat akurasi sebesar 93,67%. Algoritma *naive Bayes* mengandalkan prinsip probabilitas dalam memanfaatkan data petunjuk untuk mendukung proses pengklasifikasian keputusan.

REFERENSI

- [1] F. Ratnawati, "Implementasi Algoritma Naive Bayes Terhadap Analisis Sentimen Opini Film Pada Twitter," *INOVTEK Polbeng - Seri Inform.*, vol. 3, no. 1, p. 50, 2018, doi: 10.35314/isi.v3i1.335.
- [2] A. Sentimen *et al.*, "Analisis Sentimen Objek Wisata Bali Di Google Maps Menggunakan Algoritma Naive Bayes," *J. Sains Komput. Inform. (J-SAKTI)*, vol. 6, no. 1, pp. 418–427, 2022.
- [3] E. Febriyani and H. Februariyanti, "Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Naive Bayes Di Twitter," *J. TEKNO KOMPAK*, vol. 17, no. 2, pp. 25–38, 2022.
- [4] R. M. Cahyudi and E. B. Setiawan, "Ekspansi Fitur dengan Word2vec dalam klasifikasi Hoax di Twitter," *e-Proceeding Eng.*, vol. 10, no. 2, pp. 1765–1776, 2023.
- [5] J. J. Aripin, "Penerapan Algoritma Naive Bayes Untuk Mengklasifikasi Data Nasabah Asuransi pada BPR Pantura," 2019.
- [6] S. Lestari and S. Saepudin, "Analisis Sentimen Vaksin Sinovac Pada Twitter Menggunakan Algoritma Naive Bayes," *SISMATIK (Seminar Nas. Sist. Inf. dan Manaj. Inform.)*, pp. 163–170, 2021.
- [7] D. G. Nugroho, Y. H. Chrisnanto, and A. Wahana, "Analisis Sentimen Pada Jasa Ojek Online ... (Nugroho dkk.)," pp. 156–161, 2015.
- [8] F. Kurniawan and Q. Al Qorni, "Exploring Sentimen Analysis Using Machine Learning: A Case Study on Partai Demokrasi Indonesia Perjuangan (PDIP) in the 2024 General Election," vol. 2, no. 4, pp. 911–920, 2024.
- [9] E. Purnaningrum and I. Ariqoh, "Google Trends Analytics Dalam Bidang Pariwisata," *Maj. Ekon.*, vol. 24, no. 2, pp. 232–243, 2019, doi: 10.36456/majeko.vol24.no2.a2069.
- [10] E. Suryati, A. Ari Aldino, N. Penulis Korespondensi, and E. Suryati Submitted, "Analisis Sentimen Transportasi Online Menggunakan Ekstraksi Fitur Model Word2vec Text Embedding Dan Algoritma Support Vector Machine (SVM)," *J. Teknol. dan Sist. Inf.*, vol. 4, no. 1, pp. 96–106, 2023.
- [11] S. Sahar, "Analisis Perbandingan Metode K-Nearest Neighbor dan Naïve Bayes Clasiffier Pada Dataset Penyakit Jantung," *Indones. J. Data Sci.*, vol. 1, no. 3, pp. 79–86, 2020, doi: 10.33096/ijodas.v1i3.20.