



Analisis Sentimen Opini Twitter (X) terhadap Penggunaan ChatGPT dalam Pendidikan Tinggi dengan Perbandingan Algoritma Naïve Bayes, Support Vector Machine, dan K-Nearest Neighbors

Salma Difa Aristawidya¹, Tabita Nurtirta Purwita Sari², Hanun Syahidah Ulfya³, Adinda Aghnia Fatihin⁴, Hanif Zaidan Sinaga^{5*}

^{1,2,3,4,5}Program Studi Bisnis Digital, Fakultas Ekonomi dan Bisnis, Universitas Ibn Khaldun Bogor

¹saldifarista07@gmail.com, ²tabitapurwita@gmail.com, ³hanunsyahidahulfya04@gmail.com, ⁴adindaaghnia@gmail.com,

^{5*}hanif.zaidan@gmail.com

Abstrak

Penelitian ini bertujuan untuk menganalisis sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi serta membandingkan kinerja algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) dalam proses klasifikasi sentimen. Dataset yang digunakan diperoleh dari Kaggle dan terdiri atas 1.153 data opini yang telah memiliki label sentimen positif, netral, dan negatif. Tahapan penelitian meliputi *preprocessing* data, ekstraksi fitur menggunakan Term Frequency-Inverse Document Frequency (TF-IDF), seleksi fitur menggunakan Chi-Square, klasifikasi sentimen, serta evaluasi model menggunakan metode Stratified 10-Fold Cross Validation dengan metrik Accuracy, Precision, Recall, F1-Score, dan Confusion Matrix. Hasil penelitian menunjukkan bahwa distribusi sentimen didominasi oleh sentimen negatif sebesar 40%, diikuti sentimen positif sebesar 35%, dan sentimen netral sebesar 25%. Berdasarkan hasil perbandingan algoritma, Support Vector Machine (SVM) memperoleh performa terbaik dengan tingkat akurasi sebesar 61,75%, diikuti oleh Naïve Bayes sebesar 58,20%, dan K-Nearest Neighbors (KNN) sebesar 44,41%. Hasil tersebut menunjukkan bahwa SVM memiliki performa relatif terbaik dibandingkan Naïve Bayes dan KNN, meskipun nilai akurasi yang diperoleh menunjukkan bahwa klasifikasi opini Twitter (X) masih menghadapi tantangan akibat kompleksitas bahasa informal, ambiguitas sentimen, dan konteks penggunaan ChatGPT dalam pendidikan tinggi.

Kata Kunci: Analisis Sentimen, Chatgpt, Pendidikan Tinggi, Naïve Bayes, Support Vector Machine, K-Nearest Neighbors

1. Pendahuluan

Perkembangan Artificial Intelligence (AI) dalam beberapa tahun terakhir telah membawa perubahan yang signifikan pada berbagai sektor kehidupan, termasuk pendidikan, bisnis, dan layanan masyarakat. Salah satu teknologi AI yang berkembang pesat adalah ChatGPT yang dikembangkan oleh OpenAI. ChatGPT merupakan model bahasa berbasis Natural Language Processing (NLP) yang mampu memahami dan menghasilkan teks menyerupai bahasa manusia sehingga dapat dimanfaatkan dalam berbagai aktivitas, seperti pencarian informasi, penerjemahan bahasa, penyusunan dokumen, hingga pembuatan kode program (Wilie, 2023). Kemampuan tersebut menjadikan ChatGPT sebagai salah satu teknologi AI generatif yang banyak digunakan oleh masyarakat dan terus mengalami peningkatan pemanfaatan di berbagai bidang.

Dalam bidang pendidikan, khususnya pendidikan tinggi, ChatGPT mulai dimanfaatkan sebagai alat bantu pembelajaran yang mampu mendukung mahasiswa dalam mencari referensi, memahami materi perkuliahan, menyusun tugas, serta mengembangkan ide secara lebih efisien (Febrianty et al., 2025). Pemanfaatan teknologi ini juga dinilai mampu meningkatkan efektivitas pembelajaran dan memberikan kemudahan akses terhadap berbagai sumber informasi akademik (Pratiwi et al., 2024). Selain itu, mahasiswa cenderung memiliki persepsi yang positif terhadap penggunaan ChatGPT karena dianggap dapat membantu proses belajar, meningkatkan produktivitas, serta mendukung penyelesaian berbagai aktivitas akademik (Mariyadi et al., 2024).

Di sisi lain, penggunaan ChatGPT dalam pendidikan tinggi juga menimbulkan berbagai tantangan. Pemanfaatan yang berlebihan berpotensi menimbulkan ketergantungan terhadap teknologi sehingga dapat memengaruhi kemandirian belajar mahasiswa (Endraswari et al., 2025). Selain itu, penggunaan ChatGPT juga memunculkan kekhawatiran terkait integritas akademik, plagiarisme, serta menurunnya kemampuan berpikir kritis apabila mahasiswa terlalu bergantung pada jawaban yang dihasilkan oleh sistem AI (Supriyono et al., 2024). Pemanfaatan ChatGPT dalam lingkungan perguruan tinggi juga perlu memperhatikan aspek etika, privasi data, serta pengembangan kompetensi mahasiswa agar teknologi tersebut dapat digunakan secara bertanggung jawab (Marlin et al., 2023). Penggunaan teknologi Generative AI yang tidak terkontrol bahkan berpotensi meningkatkan risiko kecurangan akademik, disinformasi, serta ketergantungan terhadap teknologi dalam proses pembelajaran (Dewantara & Dewi, 2025). Penggunaan ChatGPT yang tidak terkontrol juga berpotensi memengaruhi

keterampilan, kolaborasi, dan kreativitas mahasiswa dalam proses pembelajaran (Putri et al., 2024). Perbedaan pandangan mengenai manfaat dan risiko tersebut menunjukkan bahwa penggunaan ChatGPT dalam pendidikan tinggi masih menjadi isu yang relevan untuk dikaji lebih lanjut.

Berbagai opini mengenai penggunaan ChatGPT banyak disampaikan melalui media sosial. Salah satu platform yang sering digunakan masyarakat untuk menyampaikan pendapat secara terbuka adalah Twitter (X). Platform ini memungkinkan pengguna untuk berbagi pengalaman, persepsi, kritik, maupun dukungan terhadap suatu fenomena secara *real-time* sehingga menjadi sumber data yang potensial untuk memahami respons publik terhadap berbagai isu, termasuk perkembangan teknologi kecerdasan buatan (Sari et al., 2024). Selain memiliki jumlah pengguna yang besar, Twitter (X) juga menjadi salah satu media sosial yang paling banyak dimanfaatkan dalam penelitian analisis sentimen karena menyediakan data opini yang beragam dan terus berkembang (Akbar et al., 2024).

Besarnya volume data opini yang dihasilkan pengguna media sosial menyebabkan proses identifikasi pandangan masyarakat secara manual menjadi kurang efektif. Oleh karena itu, diperlukan pendekatan yang mampu mengolah data teks secara otomatis dan sistematis. Salah satu pendekatan yang banyak digunakan adalah analisis sentimen. Analisis sentimen merupakan bagian dari Natural Language Processing yang digunakan untuk mengidentifikasi, mengelompokkan, dan mengevaluasi opini berdasarkan polaritas sentimen sehingga dapat memberikan gambaran mengenai kecenderungan persepsi masyarakat terhadap suatu topik secara lebih objektif dan terukur (Akbar et al., 2024).

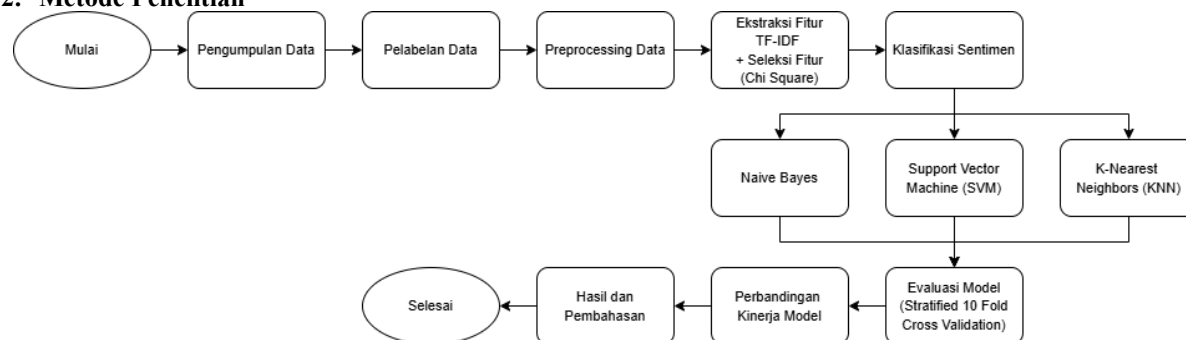
Penelitian mengenai sentimen terhadap ChatGPT telah dilakukan oleh berbagai peneliti dengan menggunakan beragam algoritma klasifikasi. Hasil penelitian menunjukkan bahwa algoritma *machine learning* mampu digunakan secara efektif untuk mengidentifikasi kecenderungan opini masyarakat terhadap ChatGPT. Wilie (2023) menunjukkan bahwa algoritma Naïve Bayes mampu digunakan untuk mengklasifikasikan opini publik terhadap ChatGPT secara efektif. Hasil serupa juga ditemukan oleh Sari et al. (2024) yang memanfaatkan algoritma Naïve Bayes untuk mengklasifikasikan opini publik terhadap ChatGPT pada platform X dengan performa klasifikasi yang sangat baik. Selain itu, Akbar et al. (2024) menggunakan algoritma K-Nearest Neighbors (KNN) dalam analisis sentimen terkait ChatGPT dan memperoleh tingkat akurasi yang tinggi. Sementara itu, Ariza et al. (2026) menunjukkan bahwa pendekatan TF-IDF dan Support Vector Machine (SVM) mampu menghasilkan performa yang baik dalam klasifikasi sentimen publik terhadap ChatGPT. Pada konteks pendidikan, Ramaputra dan Purnomo (2024) menemukan bahwa opini masyarakat terhadap penggunaan ChatGPT di bidang pendidikan memiliki kecenderungan sentimen yang beragam berdasarkan data Twitter.

Selain penelitian yang berfokus pada analisis sentimen, sejumlah penelitian juga mengkaji pemanfaatan ChatGPT dalam konteks pendidikan tinggi. Hasil penelitian menunjukkan bahwa teknologi ini memberikan berbagai manfaat dalam mendukung proses pembelajaran, namun tetap memerlukan pengawasan dalam penggunaannya. Pratiwi et al. (2024) menyatakan bahwa ChatGPT memiliki peluang untuk meningkatkan efektivitas pembelajaran, tetapi tetap memerlukan pengawasan agar tidak mengurangi kualitas proses belajar mahasiswa. Darmawan et al. (2024) menunjukkan bahwa mahasiswa memiliki persepsi dan preferensi yang positif terhadap penggunaan ChatGPT dalam proses pembelajaran. Penelitian lain juga menunjukkan bahwa penggunaan ChatGPT dapat memberikan dampak positif terhadap keterampilan, kolaborasi, dan kreativitas mahasiswa apabila dimanfaatkan secara tepat (Putri et al., 2024). Temuan-temuan tersebut menunjukkan bahwa penggunaan ChatGPT dalam pendidikan tinggi masih menjadi topik yang berkembang dan memerlukan kajian lebih lanjut dari berbagai perspektif.

Meskipun berbagai penelitian mengenai ChatGPT telah banyak dilakukan, sebagian besar penelitian terdahulu masih berfokus pada persepsi masyarakat terhadap ChatGPT secara umum atau hanya menggunakan satu algoritma klasifikasi dalam proses analisis. Penelitian yang secara khusus mengkaji opini masyarakat mengenai penggunaan ChatGPT dalam konteks pendidikan tinggi berdasarkan data Twitter (X) berbahasa Indonesia masih relatif terbatas. Selain itu, penelitian yang membandingkan kinerja algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) pada konteks penggunaan ChatGPT dalam pendidikan tinggi juga belum banyak ditemukan. Padahal, pendidikan tinggi merupakan salah satu sektor yang mengalami dampak signifikan akibat perkembangan AI generatif karena teknologi tersebut digunakan secara langsung dalam proses pembelajaran, penyusunan tugas, pencarian referensi, dan berbagai aktivitas akademik lainnya.

Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk menganalisis sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi serta membandingkan kinerja algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) dalam proses klasifikasi sentimen. Hasil penelitian diharapkan dapat memberikan gambaran mengenai kecenderungan sentimen masyarakat terhadap pemanfaatan ChatGPT dalam pendidikan tinggi serta menjadi referensi dalam menentukan algoritma klasifikasi yang lebih efektif untuk analisis sentimen berbasis data teks berbahasa Indonesia.

2. Metode Penelitian



Gambar 1. Diagram Alir Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode analisis sentimen berbasis Natural Language Processing (NLP). Penelitian dilakukan untuk mengidentifikasi kecenderungan sentimen opini pengguna Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi serta membandingkan kinerja algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) dalam proses klasifikasi sentimen. Seluruh proses pengolahan data dilakukan menggunakan Google Colab dengan memanfaatkan bahasa pemrograman Python dan pustaka pendukung *machine learning*.

2.1. Dataset Penelitian

Dataset yang digunakan dalam penelitian ini diperoleh dari platform Kaggle dengan nama dataset “Sentimen ChatGPT Mahasiswa” yang tersedia pada alamat <https://www.kaggle.com/datasets/bintangaprianta/sentimen-chatgpt-mahasiswa> dan diakses pada tanggal 4 Juni 2026. Dataset berisi 1.153 data tweet berbahasa Indonesia yang membahas penggunaan ChatGPT dalam lingkungan pendidikan tinggi.

Atribut yang digunakan dalam penelitian meliputi isi tweet sebagai data utama analisis dan label sentimen sebagai target klasifikasi. Dataset telah memiliki tiga kategori sentimen, yaitu positif (1), netral (0), dan negatif (-1). Dataset ini dipilih karena memuat opini pengguna Twitter (X) yang berkaitan dengan penggunaan ChatGPT oleh mahasiswa dalam aktivitas akademik, seperti proses pembelajaran, penyelesaian tugas, dan pencarian referensi, sehingga sesuai dengan fokus penelitian mengenai penggunaan ChatGPT dalam pendidikan tinggi.

2.2. Preprocessing Data

Tahap *preprocessing* dilakukan untuk membersihkan dan menyiapkan data teks agar dapat diproses oleh algoritma klasifikasi. Proses *preprocessing* dilakukan melalui beberapa tahapan sebagai berikut:

2.2.1. Cleaning

Menghapus karakter yang tidak diperlukan seperti URL, mention (@), hashtag (#), angka, tanda baca, emoji, dan karakter khusus lainnya.

2.2.2. Case Folding

Mengubah seluruh huruf menjadi huruf kecil (*lowercase*) untuk menyeragamkan bentuk teks.

2.2.3. Tokenization

Memecah kalimat menjadi token atau kata-kata tunggal yang akan digunakan dalam proses analisis.

2.2.4. Normalization

Mengubah kata-kata tidak baku, singkatan, atau bahasa informal ke dalam bentuk kata baku agar memiliki makna yang seragam dan lebih mudah diproses pada tahap klasifikasi.

2.2.5. Stopword Removal

Menghapus kata-kata umum yang tidak memiliki pengaruh signifikan terhadap makna sentimen, seperti “dan”, “yang”, “di”, “ke”, serta kata umum lainnya.

2.2.6. Stemming

Mengubah kata berimbuhan menjadi kata dasar sehingga variasi kata dapat diminimalkan dan proses klasifikasi menjadi lebih efektif.

2.3. Ekstraksi Fitur TF-IDF

Setelah tahap preprocessing selesai dilakukan, data teks diubah menjadi representasi numerik menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF). Metode TF-IDF digunakan untuk memberikan bobot pada setiap kata berdasarkan tingkat kemunculannya dalam suatu dokumen dan keseluruhan dokumen pada dataset. Metode ini banyak digunakan dalam penelitian analisis sentimen karena mampu merepresentasikan data teks ke dalam bentuk numerik yang dapat diproses oleh algoritma machine learning (Wilie, 2023).

Selanjutnya, dilakukan seleksi fitur menggunakan metode Chi-Square untuk memilih fitur yang paling relevan terhadap label sentimen. Seleksi fitur dilakukan untuk mengurangi fitur yang kurang informatif sehingga proses klasifikasi menjadi lebih efisien dan terarah (Ariza et al., 2026). Fitur yang terpilih kemudian digunakan sebagai

masukannya bagi algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) pada tahap klasifikasi.

2.4. Klasifikasi Sentimen

Tahap klasifikasi dilakukan dengan membandingkan tiga algoritma machine learning, yaitu Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN).

2.4.1. Naïve Bayes

Naïve Bayes merupakan algoritma klasifikasi berbasis probabilitas yang menggunakan Teorema Bayes untuk menentukan kemungkinan suatu data termasuk ke dalam kelas sentimen tertentu berdasarkan fitur yang dimiliki.

2.4.2. Support Vector Machine (SVM)

Support Vector Machine (SVM) merupakan algoritma klasifikasi yang bekerja dengan mencari *hyperplane* optimal untuk memisahkan data ke dalam kelas yang berbeda. Algoritma ini dikenal memiliki performa yang baik dalam klasifikasi data teks karena mampu menangani data berdimensi tinggi secara efektif.

2.4.3. K-Nearest Neighbors (KNN)

K-Nearest Neighbors (KNN) merupakan algoritma klasifikasi yang bekerja berdasarkan kedekatan jarak antar data. Klasifikasi dilakukan dengan menentukan kelas suatu data berdasarkan mayoritas kelas dari sejumlah tetangga terdekat yang dimilikinya.

2.5. Evaluasi Model

Evaluasi model dilakukan menggunakan metode Stratified 10-Fold Cross Validation. Metode ini membagi dataset ke dalam sepuluh bagian (*fold*) dengan mempertahankan proporsi setiap kelas sentimen pada masing-masing *fold*. Pada setiap iterasi, satu *fold* digunakan sebagai data pengujian sementara sembilan *fold* lainnya digunakan sebagai data pelatihan. Proses ini dilakukan secara bergantian hingga seluruh *fold* memperoleh kesempatan menjadi data pengujian. Penggunaan Stratified 10-Fold Cross Validation bertujuan untuk memperoleh hasil evaluasi yang lebih objektif, mengurangi bias akibat pembagian data, serta memastikan distribusi kelas sentimen tetap seimbang pada setiap proses pelatihan dan pengujian (Sari et al., 2024). Kinerja masing-masing algoritma kemudian dievaluasi menggunakan metrik Accuracy, Precision, Recall, F1-Score, dan Confusion Matrix. Selanjutnya, hasil evaluasi dari algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) dibandingkan untuk mengetahui algoritma yang menunjukkan performa relatif terbaik dalam klasifikasi sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi.

2.6. Visualisasi Hasil

Tahap akhir penelitian dilakukan dengan menyajikan hasil analisis dalam bentuk visualisasi data. Visualisasi yang digunakan meliputi distribusi sentimen, word cloud sentimen positif, netral, dan negatif, grafik perbandingan akurasi, serta confusion matrix dari masing-masing algoritma. Visualisasi tersebut digunakan untuk mempermudah penyajian dan pemahaman hasil penelitian serta mendukung analisis terhadap kecenderungan sentimen dan performa algoritma klasifikasi yang digunakan.

3. Hasil dan Pembahasan

3.1. Persiapan Data

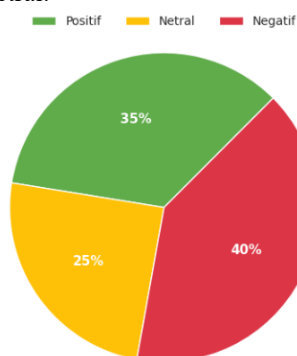
Tahap persiapan data dilakukan untuk memastikan bahwa dataset yang digunakan dalam penelitian berada dalam kondisi yang baik dan siap untuk diproses pada tahapan selanjutnya. Proses ini meliputi pengecekan data kosong (*missing values*), identifikasi data duplikat, serta analisis distribusi label sentimen yang terdapat pada dataset. Berdasarkan hasil pengecekan, tidak ditemukan data kosong pada atribut isi tweet maupun label. Selain itu, proses pemeriksaan data duplikat juga menunjukkan bahwa tidak terdapat data yang terduplikasi. Jumlah data sebelum dan sesudah proses pembersihan tetap sebanyak 1.153 data, sehingga seluruh data dapat digunakan dalam proses analisis. Selanjutnya dilakukan analisis distribusi label sentimen untuk mengetahui komposisi data yang digunakan dalam penelitian. Dataset telah memiliki tiga kategori sentimen, yaitu positif (1), netral (0), dan negatif (-1). Hasil distribusi label sentimen dapat dilihat pada Tabel 1.

Tabel 1. Distribusi Label Sentimen Dataset

Label Sentimen	Jumlah Data	Persentase
Positif (1)	403	35%
Netral (0)	285	25%
Negatif (-1)	465	40%
Total	1.153	100%

Berdasarkan Tabel 1, sentimen negatif merupakan kategori yang paling dominan dengan jumlah 465 data atau sebesar 40% dari keseluruhan dataset. Sentimen positif menempati urutan kedua dengan jumlah 403 data atau sebesar 35%, sedangkan sentimen netral memiliki jumlah paling sedikit yaitu 285 data atau sebesar 25%. Distribusi sentimen tersebut juga dapat dilihat pada Gambar 2. Grafik menunjukkan bahwa opini pengguna Twitter (X)

mengenai penggunaan ChatGPT dalam pendidikan tinggi cenderung didominasi oleh sentimen negatif dibandingkan sentimen positif maupun netral.



Gambar 2. Grafik Distribusi Label Sentimen Dataset

3.2. Hasil Preprocessing

Tahapan preprocessing yang dilakukan meliputi *cleaning*, *case folding*, *tokenization*, *normalization*, *stopword removal*, dan *stemming*. Pada tahap *stemming* digunakan pustaka Sastrawi untuk mengubah kata berimbuhan menjadi bentuk kata dasar. Hasil *preprocessing* menunjukkan bahwa teks berhasil dibersihkan dari berbagai karakter yang tidak diperlukan seperti tanda baca, simbol, dan kata-kata yang tidak memiliki kontribusi signifikan terhadap proses analisis. Selain itu, seluruh huruf berhasil dikonversi menjadi huruf kecil dan sebagian besar kata tidak baku telah dinormalisasi sehingga menghasilkan bentuk teks yang lebih seragam.

Tabel 2. Contoh Hasil Preprocessing Data

Sebelum Preprocessing	Setelah Preprocessing
Mahasiswa yang makin manja dengan ChatGPT bikin resah	mahasiswa makin manja chatgpt bikin resah
Banyak komplain bermunculan di antara para mahasiswa	banyak komplain muncul mahasiswa
Iseng nanya ChatGPT mahasiswa kesmas bisa magang	iseng tanya chatgpt mahasiswa kesmas magang

3.3. Hasil Ekstaksi Fitur TF-IDF

Setelah tahap preprocessing selesai dilakukan, data teks diubah menjadi representasi numerik menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF). Metode ini digunakan untuk memberikan bobot pada setiap kata berdasarkan tingkat kemunculannya dalam suatu dokumen dan keseluruhan dokumen pada dataset. Hasil transformasi TF-IDF menghasilkan sejumlah fitur yang merepresentasikan karakteristik teks dalam bentuk nilai numerik. Untuk meningkatkan kualitas fitur yang digunakan pada proses klasifikasi, dilakukan seleksi fitur menggunakan metode Chi-Square. Metode ini bertujuan untuk memilih fitur yang memiliki hubungan paling kuat terhadap label sentimen sehingga dapat mengurangi dimensi data dan meningkatkan efisiensi proses klasifikasi. Hasil seleksi fitur menunjukkan 10 fitur dengan nilai Chi-Square tertinggi sebagaimana ditunjukkan pada Tabel 3.

Tabel 3. Top 10 Fitur Terbaik Hasil Seleksi Chi-Square

No	Fitur	Nilai Chi-Square
1	sarjana	10.143290
2	bantu	10.034535
3	kriteria	7.209664
4	mu	5.744563
5	pro	5.596201
6	vs	5.048168
7	down	4.485207
8	nyontek	4.187101
9	in	3.986860
10	kak	3.938394

Berdasarkan Tabel 3, fitur sarjana dan bantu memiliki nilai Chi-Square tertinggi, yang menunjukkan bahwa kedua kata tersebut memiliki hubungan yang lebih kuat terhadap proses klasifikasi sentimen dibandingkan fitur lainnya. Meskipun demikian, hasil seleksi fitur masih memuat beberapa kata yang kurang informatif, seperti mu, vs, in,

dan kak. Kondisi ini menunjukkan bahwa karakteristik bahasa informal pada Twitter (X) masih menjadi salah satu keterbatasan pada tahap preprocessing, sehingga proses normalisasi dan *stopword removal* belum sepenuhnya mampu menghilangkan seluruh kata yang kurang relevan. Selanjutnya, hasil transformasi TF-IDF direpresentasikan dalam bentuk matriks fitur. Setiap nilai pada matriks menunjukkan bobot suatu kata dalam dokumen tertentu. Contoh hasil matriks TF-IDF untuk lima data pertama dapat dilihat pada Gambar 3.

	Tweet_Bersih	sarjana	bantu	kriteria	mu	pro	vs	down	nyontek	in	kak	zonauang	pedi	psikolog	gpt	presentasi
0	info chatgpt prem perplexity prem murce zonaja...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.235034	0.0	0.0	0.0	0.0
1	sempet isengisengan kalo anggar proyek mbg t p...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
2	chatgpt sebut judul judul tguas stdui kasus ma...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
3	mahasiswa makin manja chatgpt bikin resah fund...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0
4	iseng nanya chatgpt mahasiswa kesmas magang ja...	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0

Gambar 3. Matriks TF-IDF pada Lima Data Pertama

3.4. Hasil Evaluasi Model

Tahap evaluasi dilakukan untuk mengukur performa algoritma Naïve Bayes, Support Vector Machine (SVM), dan K-Nearest Neighbors (KNN) dalam mengklasifikasikan sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi. Evaluasi dilakukan menggunakan metode Stratified 10-Fold Cross Validation dengan metrik Accuracy, Precision, Recall, F1-Score, dan Confusion Matrix. Ringkasan hasil evaluasi masing-masing algoritma dapat dilihat pada Tabel 4.

Tabel 4. Ringkasan Hasil Evaluasi Kinerja Algoritma

Algoritma	Accuracy%	Precision	Recall	F1-Score
Naïve Bayes	58,20	0,66	0,58	0,51
SVM	61,75	0,63	0,62	0,58
KNN	44,41	0,46	0,44	0,42

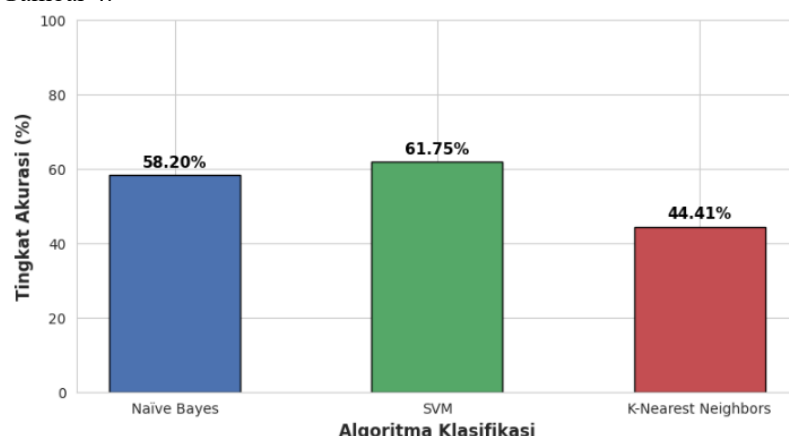
Berdasarkan Tabel 4, algoritma Support Vector Machine (SVM) memperoleh performa terbaik dengan tingkat akurasi sebesar 61,75%, diikuti oleh Naïve Bayes sebesar 58,20%, sedangkan K-Nearest Neighbors (KNN) memperoleh akurasi terendah sebesar 44,41%. Selain memiliki tingkat akurasi tertinggi, SVM juga menghasilkan nilai recall sebesar 0,62 dan F1-score sebesar 0,58, yang menunjukkan kemampuan yang lebih baik dalam mengidentifikasi dan mengklasifikasikan sentimen dibandingkan algoritma lainnya. Untuk mengetahui kinerja masing-masing algoritma pada setiap kategori sentimen, dilakukan evaluasi per kelas menggunakan metrik precision, recall, dan F1-score. Hasil evaluasi tersebut dapat dilihat pada Tabel 5.

Tabel 5. Detail Kinerja Per Kelas Setiap Algoritma

Algoritma	Sentimen	Precision	Recall	F1-Score
Naïve Bayes	Negatif	0,53	0,94	0,68
	Netral	0,84	0,06	0,11
	Positif	0,69	0,54	0,61
SVM	Negatif	0,58	0,87	0,70
	Netral	0,63	0,16	0,26
	Positif	0,68	0,65	0,66
KNN	Negatif	0,48	0,71	0,58
	Netral	0,28	0,28	0,28
	Positif	0,57	0,25	0,34

Berdasarkan Tabel 5, algoritma SVM menunjukkan performa yang paling seimbang pada ketiga kategori sentimen. Nilai F1-score tertinggi diperoleh SVM pada kelas sentimen negatif sebesar 0,70 dan sentimen positif sebesar 0,66, yang menunjukkan kemampuan yang lebih baik dalam membedakan pola sentimen dibandingkan algoritma lainnya. Sementara itu, Naïve Bayes menunjukkan performa yang cukup baik pada kelas sentimen negatif, tetapi masih mengalami kesulitan dalam mengenali sentimen netral dengan nilai recall sebesar 0,06. Kondisi ini mengindikasikan bahwa karakteristik sentimen netral memiliki kemiripan dengan sentimen positif maupun negatif sehingga lebih sulit diklasifikasikan secara tepat. Di sisi lain, KNN menghasilkan nilai precision, recall, dan F1-score yang lebih rendah pada sebagian besar kategori sentimen sehingga memperoleh performa keseluruhan yang paling rendah. Hasil tersebut menunjukkan bahwa pendekatan berbasis kedekatan jarak pada KNN kurang optimal untuk menangani data teks berdimensi tinggi yang dihasilkan dari representasi fitur TF-IDF. Secara keseluruhan, hasil evaluasi menunjukkan bahwa SVM menghasilkan performa yang lebih baik dibandingkan Naïve Bayes dan

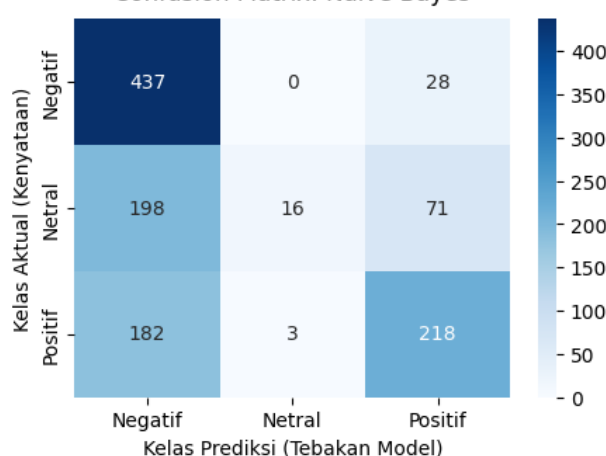
KNN pada sebagian besar metrik yang digunakan. Perbandingan tingkat akurasi ketiga algoritma secara visual dapat dilihat pada Gambar 4.



Gambar 4. Grafik Perbandingan Akurasi Algoritma

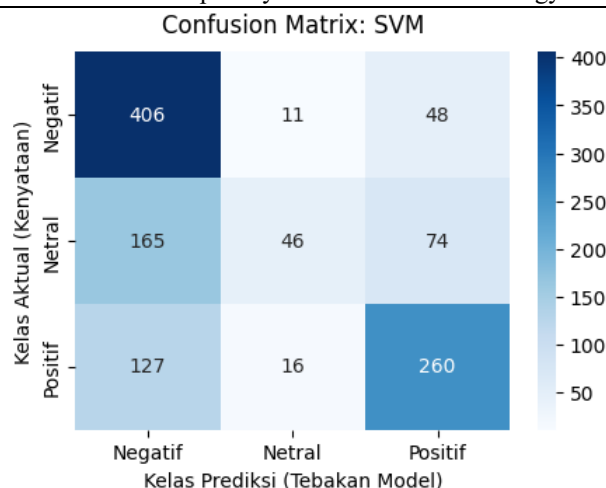
Berdasarkan Gambar 4, terlihat bahwa algoritma Support Vector Machine (SVM) menghasilkan tingkat akurasi tertinggi dibandingkan Naive Bayes dan K-Nearest Neighbors (KNN). Perbedaan tersebut menunjukkan bahwa SVM memiliki performa relatif lebih baik dalam mengklasifikasikan sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi dibandingkan dua algoritma lainnya. Sementara itu, Naive Bayes menunjukkan selisih akurasi yang relatif kecil dibandingkan SVM, sedangkan KNN menghasilkan tingkat akurasi terendah. Secara visual, grafik ini memperkuat hasil evaluasi pada Tabel 4 dan Tabel 5 yang menunjukkan adanya perbedaan performa di antara ketiga algoritma. Selain menggunakan metrik evaluasi, performa masing-masing algoritma juga dianalisis menggunakan confusion matrix untuk mengetahui kemampuan model dalam mengklasifikasikan setiap kategori sentimen secara lebih rinci. Confusion matrix pada penelitian ini diperoleh dari hasil agregasi prediksi pada seluruh *fold* selama proses evaluasi menggunakan metode Stratified 10-Fold Cross Validation. Oleh karena itu, nilai pada setiap sel confusion matrix merepresentasikan akumulasi hasil klasifikasi dari sepuluh iterasi pengujian, bukan hasil dari satu kali pembagian data latih dan data uji (*train-test split*).

Confusion Matrix: Naive Bayes



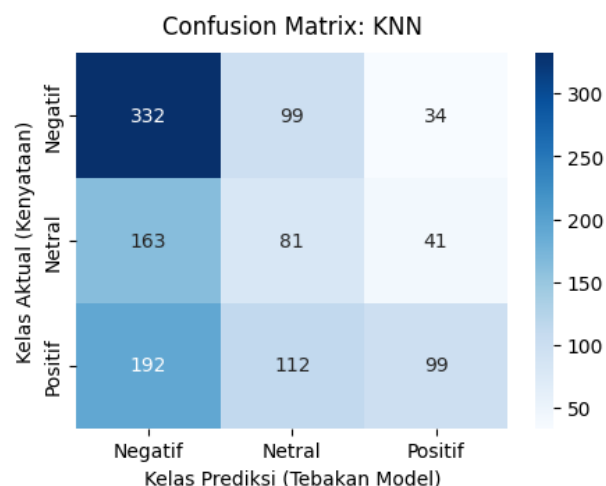
Gambar 5. Confusion Matrix Naive Bayes

Berdasarkan Gambar 5, algoritma Naive Bayes mampu mengklasifikasikan 437 data sentimen negatif dan 218 data sentimen positif dengan benar. Namun, algoritma ini masih mengalami kesulitan dalam mengenali sentimen netral, yang hanya berhasil diklasifikasikan dengan benar sebanyak 16 data dari total 285 data netral. Sebagian besar data netral diprediksi sebagai sentimen negatif sebanyak 198 data dan sentimen positif sebanyak 71 data. Hal ini menunjukkan bahwa kelas sentimen netral merupakan kategori yang paling sering mengalami kesalahan klasifikasi pada algoritma Naive Bayes. Kondisi tersebut mengindikasikan bahwa karakteristik sentimen netral memiliki kemiripan dengan dua kategori sentimen lainnya sehingga lebih sulit dibedakan oleh model.



Gambar 6. Confusion Matrix Support Vector Machine (SVM)

Berdasarkan Gambar 6, algoritma SVM menunjukkan performa yang lebih baik dibandingkan algoritma lainnya. SVM mampu mengklasifikasikan 406 data sentimen negatif, 46 data sentimen netral, dan 260 data sentimen positif dengan benar. Meskipun masih terdapat kesalahan klasifikasi pada kelas netral, jumlah prediksi yang benar pada setiap kategori sentimen relatif lebih tinggi dibandingkan Naïve Bayes dan KNN. Kelas sentimen netral tetap menjadi kategori yang paling sering salah diklasifikasikan, terutama ke dalam sentimen negatif sebanyak 165 data. Selain itu, jumlah data positif yang berhasil dikenali dengan benar juga menjadi yang tertinggi di antara ketiga algoritma. Hasil ini menunjukkan bahwa SVM lebih mampu membedakan karakteristik masing-masing kelas sentimen sehingga menghasilkan performa klasifikasi yang lebih baik dan lebih seimbang.



Gambar 7. Confusion Matrix K-Nearest Neighbors (KNN)

Berdasarkan Gambar 7, algoritma KNN menghasilkan performa yang paling rendah dibandingkan algoritma lainnya. Algoritma ini hanya mampu mengklasifikasikan 332 data sentimen negatif, 81 data sentimen netral, dan 99 data sentimen positif dengan benar. Selain itu, jumlah kesalahan klasifikasi pada setiap kategori sentimen relatif tinggi. Sebagai contoh, sebanyak 192 data sentimen positif diprediksi sebagai sentimen negatif dan 163 data sentimen netral juga diprediksi sebagai sentimen negatif. Hasil tersebut menunjukkan bahwa algoritma KNN cenderung mengelompokkan banyak data ke dalam kategori sentimen negatif sehingga menyebabkan tingginya tingkat kesalahan klasifikasi pada kelas positif dan netral. Kondisi tersebut menunjukkan bahwa pendekatan berbasis kedekatan jarak yang digunakan KNN kurang efektif dalam membedakan pola sentimen pada data teks yang memiliki jumlah fitur yang cukup banyak.

Secara keseluruhan, hasil confusion matrix menunjukkan bahwa algoritma Support Vector Machine (SVM) memiliki kemampuan klasifikasi yang paling baik dibandingkan Naïve Bayes dan K-Nearest Neighbors (KNN). SVM menghasilkan jumlah prediksi benar yang lebih tinggi pada sebagian besar kategori sentimen serta distribusi kesalahan klasifikasi yang lebih rendah. Selain itu, seluruh algoritma menunjukkan kesulitan dalam mengidentifikasi sentimen netral secara tepat, yang ditunjukkan oleh tingginya jumlah kesalahan klasifikasi pada kategori tersebut. Kondisi ini mengindikasikan bahwa karakteristik sentimen netral memiliki kemiripan dengan

3.5. Analisis Word Cloud Sentimen

3.5.1. Word Cloud Sentimen Positif



3.5.2. Word Cloud Sentimen Netral



3.5.3. Word Cloud Sentimen Negatif



Berdasarkan word cloud sentimen negatif pada Gambar 10, kata yang paling dominan adalah chatgpt, mahasiswa, aku, dosen, pakai, jadi, banget, dan kuliah. Selain itu, terdapat beberapa kata seperti nyontek, nilai, tugas, dan jawab yang juga cukup sering muncul. Dominasi kata-kata tersebut menunjukkan bahwa sentimen negatif cenderung berkaitan dengan kekhawatiran terhadap dampak penggunaan ChatGPT dalam aktivitas akademik. Hasil ini menunjukkan bahwa meskipun ChatGPT memberikan berbagai manfaat, sebagian pengguna masih memiliki persepsi negatif terhadap penggunaannya dalam lingkungan pendidikan tinggi, terutama yang berkaitan dengan aspek integritas akademik.

3.6. Pembahasan

Berdasarkan hasil evaluasi model, algoritma Support Vector Machine (SVM) memperoleh performa terbaik dengan tingkat akurasi sebesar 61,75%, diikuti oleh Naïve Bayes sebesar 58,20%, sedangkan K-Nearest Neighbors (KNN) memperoleh akurasi terendah sebesar 44,41%. Hasil ini menunjukkan bahwa SVM lebih efektif dalam mengklasifikasikan sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi dibandingkan dua algoritma lainnya. Kemampuan SVM dalam menangani data berdimensi tinggi yang dihasilkan dari representasi TF-IDF memungkinkan algoritma ini membentuk batas pemisah antar kelas sentimen dengan lebih optimal sehingga menghasilkan tingkat klasifikasi yang lebih akurat. Hasil tersebut mendukung penelitian Ariza et al. (2026) yang menunjukkan bahwa algoritma SVM memiliki performa yang baik dalam klasifikasi sentimen terkait ChatGPT.

Meskipun memiliki akurasi yang lebih rendah dibandingkan SVM, algoritma Naïve Bayes tetap menunjukkan performa yang cukup baik dengan selisih akurasi yang relatif kecil. Hal ini menunjukkan bahwa pendekatan probabilistik yang digunakan Naïve Bayes masih mampu mengidentifikasi hubungan antara fitur teks dan kategori sentimen secara efektif. Sebaliknya, algoritma K-Nearest Neighbors (KNN) menghasilkan performa terendah karena proses klasifikasinya sangat bergantung pada kedekatan jarak antar data. Pada data teks yang memiliki jumlah fitur yang banyak, proses pengukuran jarak antar data menjadi lebih kompleks sehingga memengaruhi kemampuan KNN dalam melakukan klasifikasi secara akurat. Perbedaan performa tersebut menunjukkan bahwa karakteristik data teks lebih sesuai diproses menggunakan algoritma yang mampu menangani data berdimensi tinggi secara efektif, seperti SVM.

Keberagaman sentimen yang ditemukan dalam penelitian ini menunjukkan bahwa penggunaan ChatGPT dalam pendidikan tinggi tidak hanya dipandang sebagai alat bantu pembelajaran, tetapi juga sebagai teknologi yang memunculkan berbagai pertimbangan etis dalam lingkungan akademik. Kondisi tersebut sejalan dengan penelitian Yulianti et al. (2025) yang menunjukkan bahwa persepsi masyarakat terhadap penggunaan ChatGPT dalam pendidikan dipengaruhi oleh berbagai aspek, seperti manfaat sebagai asisten belajar dan implikasinya terhadap etika akademik. Hasil ini juga memperkuat penelitian Lindiananda (2026) yang menunjukkan bahwa penggunaan ChatGPT memiliki keterkaitan dengan aktivitas belajar mahasiswa sehingga menghasilkan beragam tanggapan positif, netral, maupun negatif.

Dominasi sentimen negatif sebesar 40% menunjukkan bahwa sebagian pengguna Twitter (X) masih memiliki kekhawatiran terhadap penggunaan ChatGPT dalam pendidikan tinggi. Hasil analisis word cloud sentimen negatif menunjukkan kemunculan kata-kata seperti nyontek, tugas, nilai, dan jawab yang mengindikasikan adanya persepsi bahwa ChatGPT berpotensi digunakan untuk menyelesaikan tugas secara instan tanpa melalui proses berpikir yang memadai. Selain itu, penggunaan ChatGPT juga dikhawatirkan dapat meningkatkan ketergantungan mahasiswa terhadap teknologi sehingga berpotensi mengurangi kemandirian belajar dan kemampuan berpikir kritis. Kekhawatiran lainnya berkaitan dengan integritas akademik, terutama terkait kemungkinan penyalahgunaan ChatGPT dalam menyelesaikan tugas atau memperoleh jawaban secara cepat tanpa proses pembelajaran yang optimal. Oleh karena itu, pemanfaatan ChatGPT di lingkungan perguruan tinggi perlu diimbangi dengan kebijakan dan pedoman penggunaan yang jelas agar teknologi tersebut dapat dimanfaatkan sebagai alat bantu pembelajaran tanpa mengurangi kualitas proses belajar maupun integritas akademik mahasiswa.

4. Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa opini pengguna Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi didominasi oleh sentimen negatif sebesar 40%, diikuti sentimen positif sebesar 35%, dan sentimen netral sebesar 25%. Dominasi sentimen negatif menunjukkan bahwa meskipun ChatGPT memberikan berbagai manfaat dalam mendukung aktivitas akademik, masih terdapat sejumlah kekhawatiran terkait dampak penggunaannya dalam lingkungan pendidikan tinggi, terutama yang berkaitan dengan integritas akademik, ketergantungan terhadap teknologi, dan kualitas proses pembelajaran.

Berdasarkan hasil perbandingan algoritma, Support Vector Machine (SVM) memperoleh tingkat akurasi sebesar 61,75%, diikuti oleh Naïve Bayes sebesar 58,20%, dan K-Nearest Neighbors (KNN) sebesar 44,41%. Hasil tersebut menunjukkan bahwa SVM memiliki performa relatif terbaik dibandingkan Naïve Bayes dan KNN dalam mengklasifikasikan sentimen opini Twitter (X) terhadap penggunaan ChatGPT dalam pendidikan tinggi. Namun, tingkat akurasi yang diperoleh juga menunjukkan bahwa model masih memiliki keterbatasan dalam menangani

kompleksitas opini pengguna Twitter (X) yang beragam dan dinamis. Oleh karena itu, penelitian selanjutnya dapat mempertimbangkan penggunaan metode atau algoritma lain yang berpotensi menghasilkan performa klasifikasi yang lebih baik.

Reference

- Akbar, Y., Regita, A. N. H., Sugiyono, & Wahyudi, T. (2024). Analisa Sentimen Pada Media Sosial “ X ” Pencarian Keyword Chatgpt Menggunakan Algoritma K-Nearest Abstrak. *Jurnal Indonesia : Manajemen Informatika Dan Komunikasi*, 5(3), 3291–3305. <https://doi.org/10.35870/jimik.v5i3.1016>
- Ariza, R., Fatchan, M., & Suprianto, A. (2026). Analisis Sentimen Opini Publik Terhadap Chatgpt Pada Platform X Menggunakan Pendekatan Tf-Idf Dan Support Vector Machine. *JACIS: Journal Automation Computer Information System*, 6(1), 150–161. <https://doi.org/https://doi.org/10.47134/Jacis.V6i1.179>
- Darmawan, I. K. A., Supriyadi, & Junaidi. (2024). ANALISIS PRESEPSI DAN PREFERENSI MAHASISWA TERHADAP PENGGUNAAN CHATGPT DALAM PROSES PEMBELAJARAN (STUDI KASUS DI UNIVERSITAS SAMAWA FAKULTAS EKONOMI DAN MANAJEMEN). *Jurnal Tambora*, 8(2), 10–24. <https://doi.org/https://doi.org/10.36761/Suffix>
- Dewantara, B. A., & Dewi, L. K. (2025). Generative AI Dalam Pembelajaran Mahasiswa: Antara Inovasi Pendidikan Dan Integritas Akademik. *JIP (Jurnal Ilmiah Ilmu Pendidikan)*, 8(7), 8209–8217.
- Endraswari, P. M., Tou, N., & Zaliman, I. (2025). ANALISIS KETERGANTUNGAN MAHASISWA TERHADAP CHATGPT DALAM KONTEKS AKADEMIK: PEMODELAN KLASIFIKASI BERBASIS NAÏVE BAYES. *Jurnal Informasi Interaktif*, 10(3), 199–208.
- Febrianty, C., Sari, M. T. P., & Syarafi, R. H. (2025). ANALISIS DAMPAK CHATGPT TERHADAP PROSES PEMBELAJARAN MAHASISWA: SYSTEMATIC LITERATURE REVIEW. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 9(1), 949–961.
- Lindiananda, N. (2026). Pengaruh Intensitas Penggunaan Chatgpt Terhadap Kemandirian Belajar Mahasiswa. *Jurnal Teknologi Pendidikan*, 11(1), 41–49.
- Mariyadi, Asri, D. F. Al, Pratama, I. W. S. W. P., Marthaningrum, G. B., & Siahaan, J. G. (2024). Persepsi Mahasiswa Terhadap Pemanfaatan Chatgpt Dalam Proses Pembelajaran Di Perguruan Tinggi. *Didaktika: Jurnal Kependidikan*, 13(4), 5423–5438.
- Marlin, K., Tantrisa, E., Mardikawati, B., Anggraini, R., & Susilawati, E. (2023). Manfaat Dan Tantangan Penggunaan Artificial Intelligences (AI) Chat GPT Terhadap Proses Pendidikan Etika Dan Kompetensi Mahasiswa Di Perguruan Tinggi. *INNOVATIVE: Journal Of Social Science Research Volume*, 3(6), 5192–5201.
- Pratiwi, N. K., Yulianto, B., Mintowati, Supratno, H., Sodiq, S., & Mulyono. (2024). Persepsi Mahasiswa Terhadap Penggunaan Chatgpt : Peluang Dan Tantangan Bagi Pembelajaran Bahasa Indonesia Sebagai Mata Kuliah Wajib Pada Kurikulum Perguruan Tinggi. *Jurnal Onoma: Pendidikan, Bahasa Dan Sastra*, 10(3), 2727–2742.
- Putri, Z. H. A., Pradana, N. R., Yustraini, Y. A., & Efansyah, A. D. (2024). Analisis Pengaruh Chat GPT Terhadap Keterampilan , Kolaborasi , Dan Kreativitas Mahasiswa: Metode Systematic Literature Review Identifikasi Dampak Dan Pengaruh. *INNOVATIVE: Journal Of Social Science Research Volume*, 4(2), 7983–7999.
- Ramaputra, M. G., & Purnomo, H. (2024). Analisis Sentimen Opini Masyarakat Terhadap Penggunaan Chatgpt Di Bidang Pendidikan Berbasis Twitter. *Jurnal Pepadun*, 5(3), 275–285.
- Sari, Y. K., Rozi, F., Muhyiddin, S., & Sukmana, F. (2024). Sentiment Analysis Of Public Opinion On Application X (Twitter) In Indonesia Against Chatgpt Using Naïve Bayes Algorithm. *JUPI (Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika)*, 9(4), 2473–2484.
- Supriyono, A., Prihandono, T., & Lesmono, A. D. (2024). Dampak Dan Tantangan Pemanfaatan Chatgpt Dalam Pembelajaran Pada Kurikulum Merdeka: Tinjauan Literatur Sistematis The Impact And Challenges Of Utilizing Chatgpt In Learning Within The Kurikulum : A Systematic Literature Review. *Jurnal Pendidikan Dan Kebudayaan*, 9(2), 9–12. <https://doi.org/https://doi.org/10.24832/Jpnk.V9i2.5214>
- Wilie, D. P. (2023). Analisis Sentimen Opini Publik Terhadap Chatgpt Di Twitter Menggunakan Metode Naive Bayes. *Jurnal Nasional Ilmu Komputer*, 4(4), 35–44.
- Yulianti, D., Kristianti, N., Saragih, A. S., & Leonardo, T. (2025). Aspect-Based Sentiment Analysis Penggunaan Chatgpt Dalam Pendidikan: Perbandingan Model LSTM, Bi-LSTM, Dan CNN. *JOINTECOMS (Journal Of Information Technology And Computer Science)*, 5(4), 10–12.