

ANALISIS SENTIMEN CHATGPT DATA SOSIAL MEDIA X(TWITTER) DENGAN MENGGUNAKAN FINE TUNING XL NET

Theresia Hendrawati¹⁾, Ni Luh Wiwik Sri Rahayu Ginantra²⁾

^{1,2}Prodi Informatika, Teknologi dan Informatika, Institut Bisnis dan Teknologi Indonesia

Correspondence author: N.L.W.S.R.Ginantra, wiwik@instiki.ac.id, Denpasar, Indonesia

Abstract

ChatGPT (Generative Pre-training Transformer) is an artificial intelligence technology designed to mimic human conversation in text form and has become an important tool in various fields, including education. This study aims to analyze public sentiment toward the use of ChatGPT, which can be categorized into positive and negative sentiments. The data for the study was obtained from 5,686 user reviews on the Twitter platform, collected through Google Colaboratory and processed with pre-processing steps. The data was labeled as positive and negative, then classified using fine-tuning on the XLNet model, a Transformer-based language model. The results show that the fine-tuned XLNet model achieved an accuracy of 88.45%, a precision of 89%, a recall of 88%, and an F1-score of 89%, as measured using the Confusion Matrix. This study demonstrates that fine-tuning XLNet is effective in classifying the sentiment of ChatGPT user reviews related to education.

Keywords: chatgpt, fine tuning xl net, sentiment classification, twitter

Abstrak

ChatGPT (*Generative Pre-training Transformer*) adalah teknologi kecerdasan buatan yang dirancang untuk menirukan percakapan manusia dalam bentuk teks dan telah menjadi alat penting di berbagai bidang, termasuk pendidikan. Penelitian ini bertujuan untuk menganalisis akurasi kinerja model *fine tuning* XL Net dari sentimen masyarakat terhadap penggunaan ChatGPT, yang dapat dikategorikan menjadi sentimen positif dan negatif. Data penelitian diperoleh dari 5.686 ulasan pengguna di platform Twitter, dikumpulkan melalui Google Colaboratory dan diproses dengan tahap *pre-processing*. Data diberi label positif dan negatif, lalu diklasifikasikan menggunakan metode *fine-tuning* pada model XLNet, model bahasa berbasis *Transformer*. Hasil penelitian menunjukkan bahwa model *fine-tuning* XLNet mencapai akurasi 88,45%, precision 89%, recall 88%, dan F1-score 89%, yang diukur menggunakan Confusion Matrix. Penelitian ini membuktikan bahwa fine-tuning XLNet efektif dalam mengklasifikasikan sentimen ulasan pengguna ChatGPT terkait pendidikan.

Kata Kunci: chatgpt, *fine tuning* xl net, klasifikasi sentimen, twitter

A. PENDAHULUAN

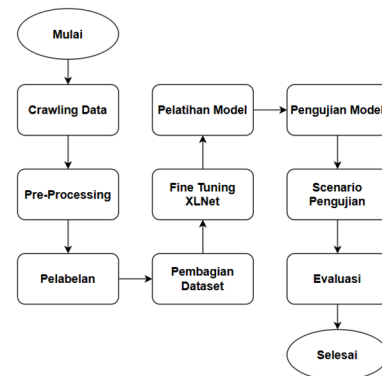
ChatGPT, yang dirilis pada 30 November 2022 oleh OpenAI, merupakan program kecerdasan buatan yang berinteraksi melalui percakapan (Singh et al., 2023). Pada Oktober 2023, chat.openai.com mencatat 180,5 juta pengguna dan 1,7 miliar kunjungan, meningkat 170% dari Februari 2023 yang sebesar satu miliar, dan sekitar enam kali lebih banyak dari dua ratus enam puluh enam juta kunjungan Desember 2022 (Taecharungroj, 2023). Meski berpotensi memajukan pendidikan, seperti meningkatkan motivasi belajar (Murcahyanto, 2023), penggunaannya juga dikhawatirkan menurunkan kemampuan berpikir kritis pelajar (Aiman & Kurniawaty, 2023). Karena itu, penting menganalisis sentimen pengguna terhadap ChatGPT, terutama di Twitter yang banyak digunakan oleh pengguna produktif di Indonesia (Pasek et al., 2022).

Penelitian sebelumnya terkait sentimen terhadap ChatGPT antara lain: analisis di Twitter dengan metode *C45* dan *Naïve Bayes* yang menghasilkan akurasi 77,33% (Akbar & Sugiharto, 2023); penggunaan *K-Nearest Neighbors* dengan akurasi 62% (Ilham & Ilham, 2023); serta kombinasi *SVM* dan *Naïve Bayes* yang mencapai akurasi 55%.

Penelitian ini menggunakan *fine-tuning XLNet* untuk menganalisis sentimen pengguna Twitter terhadap ChatGPT. *Fine-tuning* merupakan proses pelatihan tambahan untuk meningkatkan kinerja model bahasa. *XLNet* menggunakan pendekatan permutasi yang memungkinkan pemahaman konteks dua arah (Topal et al., 2021), sehingga diharapkan mampu memberikan hasil analisis sentimen yang lebih akurat.

B. METODE PENELITIAN

Agar penelitian ini dapat dilaksanakan, terdapat beberapa tahapan atau alur yang akan dilakukan. Rancangan tahapan penelitian dilihat pada Gambar 1.



Gambar 1. Rancangan Penelitian

Pengumpulan Data

Data dikumpulkan dengan metode *crawling* dari Twitter menggunakan *Python* di *Google Colab*, pengumpulan data dari November 2022 hingga Desember 2023. Pemilihan Twitter sebagai sumber data didasarkan, kemampuannya menyebarkan informasi dengan cepat dibandingkan platform lainnya (Wandani, 2021).

Pre-processing

Pre-processing adalah tahap pembersihan data mentah untuk membuatnya lebih terstruktur dan siap dianalisis. Tahapan ini meliputi penghapusan data duplikat, penghilangan simbol atau karakter khusus, serta penyaringan kata-kata yang tidak relevan (Santoso & Wibowo, 2022).

Pelabelan

Pelabelan data dilakukan untuk mengklasifikasikan data ke dalam kelas positif atau negatif menggunakan kamus *Insert Lexicon*, yang berisi kata-kata beserta bobot untuk menentukan nilai polaritas. Ulasan dikategorikan sebagai sentimen positif jika nilai polaritas lebih dari nol, dan negatif jika nilainya di bawah nol (Faz rinand Andeswari, 2022).

Pembagian Data

Pembagian data dilakukan dengan rasio 80% dan 20%. Dari 5.594 data yang dikumpulkan antara November 2022 hingga Desember 2023, 80% (3.356 data) digunakan

untuk pelatihan, dan 20% (1.118 data) untuk pengujian.

Klasifikasi Data

Klasifikasi data bertujuan menentukan kategori sentimen tweet setelah proses *pre-processing* dan pelabelan. Penelitian ini menggunakan *Fine-Tuning XLNet*, model *deep learning* berbasis pemodelan bahasa permutasi yang dikembangkan oleh Carnegie Mellon University dan Google (Dhivyaa & Nithya, 2023). *XLNet* menggabungkan keunggulan *autoregressive* dan *autoencoding* untuk memahami konteks dua arah tanpa kehilangan informasi. *Fine-tuning* dilakukan dengan menyesuaikan model *pre-trained* menggunakan *supervised learning* serta pengaturan *hyperparameter* seperti *epoch*, *learning rate*, dan *batch size* guna mengoptimalkan performa klasifikasi sentimen terkait pendidikan dan ChatGPT (Li et al., 2019).

Evaluasi Model

Confusion matrix digunakan untuk mengevaluasi kinerja model dalam klasifikasi, dengan menganalisis efektivitas pada setiap tahap pelatihan dan pengujian. Mencakup empat metrik utama: *Recall*, *Precision*, *Accuracy*, dan *F1-Score*, yang dihitung dari *output confusion matrix* (Indrayuni, 2019).

C. HASIL DAN PEMBAHASAN

Pengumpulan data opini pengguna Twitter terkait ChatGPT dalam pendidikan, menggunakan "chatGPT_pendidikan" melalui *Google Colab* dengan *Python* dan *library Twitter crawler*.

Pada Gambar 2 menampilkan hasil *crawling* data mentah dengan kata kunci "chatGPT_pendidikan" pada periode 1 November 2022–30 November 2023, dengan total 5.594 data mentah yang selanjutnya diproses melalui tahapan *pre-processing*.

full_text	created_at	username	tweet_url
Senang lihat anak SD bisa memanfaatkan ChatGPT untuk belajar	Wed Mar 29 23:56:22 +0000 2023	farahgalli	https://twitter.com/farahgalli/status/1641222737908424448
ChatGPT na Sodiql AGI Setara Kecerdasan Manusia Apalagi ini Perkembangan pesat dalam kecerdasan buatan AI terus menjadi perhatian tahun ini sering kemunculan ChatGPT Mesin telusur yang bisa menulis chatbot AI sendiri serta ramainya pengguna AI membuat banyak orang waspada	Wed Mar 29 23:50:27 +0000 2023	gabbbz	https://twitter.com/gabbbz/status/1641226351486136920
Pemerintah Bakal Blokir ChatGPT? Pada akhir Februari 2023 muncul kabar bahwa pemerintah berencana untuk memblokir ChatGPT Alasan pemblokiran ChatGPT oleh pemerintah penyebabnya bahwa ChatGPT belum melakukan Pendaftaran Sistem Elektronik	Wed Mar 29 23:39:15 +0000 2023	ChoNunik	https://twitter.com/ChoNunik/status/164122523334256009
Sama rder Tapi aku bisanya jadi paham waktu ngerjain tugas rder Soalnya waktu nugas aku diskusi sama chatgpt Saling bertukar pemahaman kita	Wed Mar 29 22:31:13 +0000 2023	AryBandung	https://twitter.com/6704550604

Gambar 2. Hasil *Crawling* Data

Pre-Processing

Setelah memperoleh 5.594 data kotor dari proses *crawling*, tahap selanjutnya adalah *pre-processing*, dimulai dengan :

1. *Cleaning*: Tahap pertama *pre-processing* adalah *cleaning*, yang menghapus elemen tidak relevan seperti *URL*, tanda baca, *mention*, *emoticon*, *tag HTML*, dan lainnya (Musfiroh et al., 2021). Hasil *cleaning* dapat dilihat pada Gambar 3.

Full_text	Clean
"RT@narasiv Teknologi AI seperti ChatGPT benar-benar mengubah cara kita bekerja dan belajar. 🌟 #inovasi"	teknologi ai seperti chatgpt benar-benar mengubah cara kita bekerja dan belajar
"RT@suhaaa_th ChatGPT selalu siap membantu dengan jawaban cepat dan akurat. Sangat membantu! 🙌 #Teknologi"	chatgpt selalu siap membantu dengan jawaban cepat dan akurat sangat membantu
Menurutku @anisuhadah_ChatGPT lebih baik dari pada google untuk membantu saya menemukan dan memperbaiki error kalo pake google perlu milih artikel dulu tapi dengan chatGPT langsung dapat penjelasannya langkah langkahnya bahkan di kasih contoh Yang paling saya suka dia bisa bahasa Indonesia https://tco/YySwjSGH8	Menurutku ChatGPT lebih baik dari pada google untuk membantu saya menemukan dan memperbaiki error kalo pake google perlu milih artikel dulu tapi dengan chatGPT langsung dapat penjelasannya langkah langkahnya bahkan di kasih contoh Yang paling saya suka dia bisa bahasa Indonesia

Gambar 3. Hasil *Cleaning*

2. *Tokenizer*: Proses ini memecah kalimat menjadi satuan token (kata) untuk mengidentifikasi setiap kata secara terpisah dalam teks (Santoso & Wibowo, 2022). Hasil *tokenizer* dapat dilihat pada Gambar 4.

Full_text	Clean
Jujurly chatgpt ini sangat membantuku mulai dari cari topik diskusi n memperkuat argumen smpe akhirnya aku berhasil bikin skripsi n sidang dalam 6 bulan Yah tpi ofc aku gk cuma dibantu chat gpt doang aku juga pake humata les dikit soal topik tugas akhirku sma support temen dy?	['jujurly', 'chatgpt', 'membantuku', 'cari', 'topik', 'diskusi', 'n', 'memperkuat', 'argumen', 'smpe', 'berhasil', 'bikin', 'skripsi', 'n', 'sidang', '6', 'yah', 'tpi', 'ofc', 'gk', 'dibantu', 'chat', 'gpt', 'doang', 'pake', 'humata', 'les', 'dikit', 'topik', 'tugas', 'akhirku', 'sma', 'support', 'temen']
ChatGPT tidak hadir begitu saja dengan segala kecerdasannya Ia dirancang dengan beragam data yang dikumpulkan oleh para pengembang maupun perusahaan Termasuk data data toksisitas linguistik //tco/gUwMBcwnhM	['chatgpt', 'hadir', 'kecerdasannya', 'dirancang', 'beragam', 'data', 'dikumpulkan', 'pengembang', 'perusahaan', 'data data', 'toksisitas', 'linguistik']
bantu coba kritik pke chatgpt silahkan komen nya //tco/3gbVtIsORX	['bantu', 'coba', 'kritik', 'pke', 'chatgpt', 'silahkan', 'komen', 'nya']

Gambar 4. Hasil *Tokenizer*

3. *Transform Case*: Proses ini mengubah semua huruf menjadi huruf kecil untuk menyeragamkan format teks (Santoso & Wibowo, 2022). Hasil *transform case* dapat dilihat di Gambar 5.

Full text	Clean
Jujurly chatgpt ini sangat membantuku mulai dari cari topik diskusi n memperkuat argumen smpe akhirnya aku berhasil bikin skripsi n sidang dalam 6 bulan Yah tpi ofc aku gk cuma dibantu chat gpt doang aku juga pake humata les dikit soal topik tugas akhirku sma support temen 6Y?	['jujurly', 'chatgpt', 'membantuku', 'cari', 'topik', 'diskusi', 'n', 'memperkuat', 'argumen', 'smpe', 'berhasil', 'bikin', 'skripsi', 'n', 'sidang', '6', 'yah', 'tpi', 'ofc', 'gk', 'dibantu', 'chat', 'gpt', 'doang', 'pake', 'humata', 'les', 'dikit', 'topik', 'tugas', 'akhirku', 'sma', 'support', 'temen']
ChatGPT tidak hadir begitu saja dengan segala kecerdasannya la dirancang dengan beragam data yang dikumpulkan oleh para pengembang maupun perusahaan Termasuk datadata toksisitas linguistik //tco/gUwMBcwnhM bantu coba kritik pke chatgpt silahkan komen nya //tco/3gbVtIsORX	['chatgpt', 'hadir', 'kecerdasannya', 'dirancang', 'beragam', 'data', 'dikumpulkan', 'pengembang', 'perusahaan', 'datadata', 'toksisitas', 'linguistik']
	['bantu', 'coba', 'kritik', 'pke', 'chatgpt', 'silahkan', 'komen', 'nya']

Gambar 5. Hasil *Transform Case*

4. *Filter Stopwords*: Proses ini menghapus kata-kata umum seperti "dan", "yang", "di" menggunakan kamus *stopwords* bahasa Indonesia (Ratniasih and Jayanti, 2023). Hasilnya ditampilkan pada Gambar 6.

Full text	Clean	Filter_stopword
Jujurly chatgpt ini sangat membantuku mulai dari cari topik diskusi n memperkuat argumen smpe akhirnya aku berhasil bikin skripsi n sidang dalam 6 bulan Yah tpi ofc aku gk cuma dibantu chat gpt doang aku juga pake humata les dikit soal topik tugas akhirku sma support temen 6Y?	['jujurly', 'chatgpt', 'membantuku', 'cari', 'topik', 'diskusi', 'n', 'memperkuat', 'argumen', 'smpe', 'berhasil', 'bikin', 'skripsi', 'n', 'sidang', '6', 'yah', 'tpi', 'ofc', 'gk', 'dibantu', 'chat', 'gpt', 'doang', 'pake', 'humata', 'les', 'dikit', 'topik', 'tugas', 'akhirku', 'sma', 'support', 'temen']	['jujurly', 'chatgpt', 'membantuku', 'cari', 'topik', 'diskusi', 'memperkuat', 'argumen', 'smpe', 'berhasil', 'bikin', 'skripsi', 'sidang', 'dibantu', 'chat', 'gpt', 'pake', 'humata', 'les', 'topik', 'tugas', 'akhirku', 'support', 'temen']
ChatGPT tidak hadir begitu saja dengan segala kecerdasannya la dirancang dengan beragam data yang dikumpulkan oleh para pengembang maupun perusahaan Termasuk datadata toksisitas linguistik //tco/gUwMBcwnhM bantu coba kritik pke chatgpt silahkan komen nya //tco/3gbVtIsORX	['chatgpt', 'hadir', 'kecerdasannya', 'dirancang', 'beragam', 'data', 'dikumpulkan', 'pengembang', 'perusahaan', 'datadata', 'toksisitas', 'linguistik']	['chatgpt', 'kecerdasannya', 'data', 'datadata', 'linguistik']
	['bantu', 'coba', 'kritik', 'pke', 'chatgpt', 'silahkan', 'komen', 'nya']	['bantu', 'coba', 'kritik', 'chatgpt', 'komen']

Gambar 6. Hasil *Stopwords*

Pelabelan Data

Setelah *pre-processing*, data dilabeli menggunakan *Microsoft Excel* dengan metode Liu Hu: skor >0 diberi label positif, <0 negatif, dan =0 netral. Karena hanya digunakan kategori positif dan negatif, data netral dihapus, menyisakan 4.849 data untuk tahap selanjutnya. Hasil pelabelan ditampilkan pada Gambar 7.

Clean	Filter_stopword	Sentiment
['jujurly', 'chatgpt', 'membantuku', 'cari', 'topik', 'diskusi', 'n', 'memperkuat', 'argumen', 'smpe', 'berhasil', 'bikin', 'skripsi', 'n', 'sidang', '6', 'yah', 'tpi', 'ofc', 'gk', 'dibantu', 'chat', 'gpt', 'doang', 'pake', 'humata', 'les', 'dikit', 'topik', 'tugas', 'akhirku', 'sma', 'support', 'temen']	['jujurly', 'chatgpt', 'membantuku', 'cari', 'topik', 'diskusi', 'memperkuat', 'argumen', 'smpe', 'berhasil', 'bikin', 'skripsi', 'sidang', 'dibantu', 'chat', 'gpt', 'pake', 'humata', 'les', 'topik', 'tugas', 'akhirku', 'support', 'temen']	1
['chatgpt', 'hadir', 'kecerdasannya', 'dirancang', 'beragam', 'data', 'dikumpulkan', 'pengembang', 'perusahaan', 'datadata', 'toksisitas', 'linguistik']	['chatgpt', 'kecerdasannya', 'data', 'datadata', 'linguistik']	1
['bantu', 'coba', 'kritik', 'pke', 'chatgpt', 'silahkan', 'komen', 'nya']	['bantu', 'coba', 'kritik', 'chatgpt', 'komen']	1

Gambar 7. Hasil Pelabelan Data

Pembagian Data

Sebelum klasifikasi, data yang sudah dilabeli dibagi menjadi dua bagian: 3.879 data untuk *training* dan 970 data untuk test menggunakan *library Scikit-learn*. Hasil pembagian data dapat dilihat pada Gambar 8.

Data Type	Count
Total data	4849
Training data	3879
Testing data	970

Gambar 8. Pembagian Data

Fine Tuning XLNet

Penelitian ini menerapkan *fine-tuning* pada model *XLNet*, yang menggunakan metode permutasi untuk memahami konteks lebih baik dari *BERT*, dengan bantuan *library Transformers* dari *HuggingFace* yang mendukung *PyTorch* dan *TensorFlow* (Li et al., 2019).

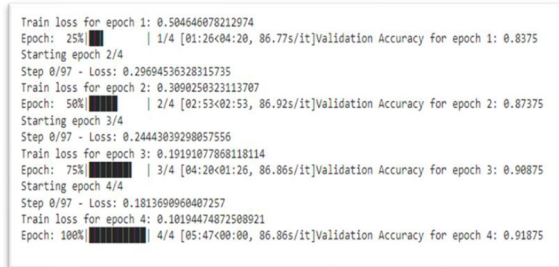
Data Preparation

Sebelum pelatihan, data diproses dengan *XLNet Tokenizer* yang mengubah teks menjadi token sesuai *vocabulary XLNet*, menambahkan token khusus seperti [SEP] dan [CLS] untuk memisahkan segmen dan mewakili kalimat dalam klasifikasi.

Tahapan Training

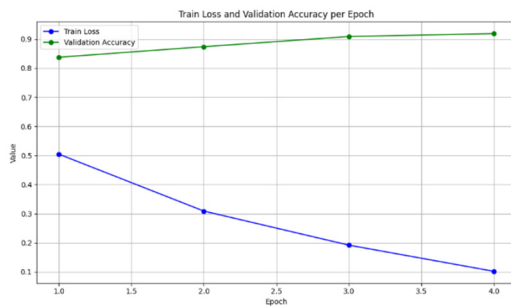
Pelatihan dilakukan dengan *fine-tuning* pada model *pre-trained XLNet* yang memiliki 12 lapisan *encoder*. Dengan *hyperparameter batch size 32, epoch 4, dan learning rate 2e-5*, model menunjukkan penurunan *train loss* dari 0,5046 (*epoch 1*) menjadi 0,1813 (*epoch*

4), sementara *validation accuracy* meningkat dari 0,8375 menjadi 0,9187. Hasil pelatihan ditampilkan pada Gambar 9.



Gambar 9. Hasil Pelatihan *XLNet*

Perubahan *train loss* yang menurun dan *validation accuracy* yang meningkat menunjukkan model semakin baik dalam mempelajari data tanpa *overfitting* signifikan. Hasil lengkap ditampilkan pada Gambar 10.



Gambar 10. Hasil *Train Loss* dan Validasi

Pengujian Model *XLNet*

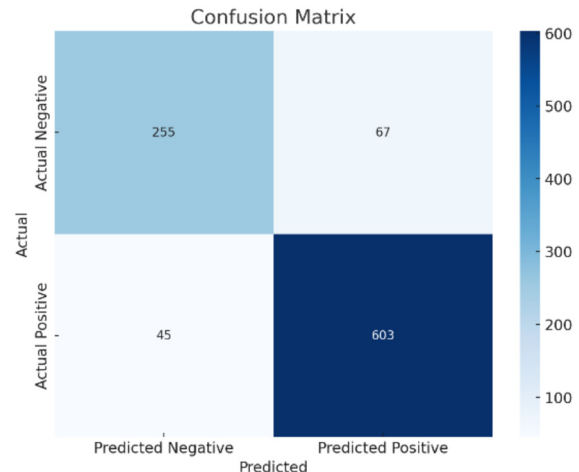
Proses ini dilakukan untuk menguji model *XLNet* menggunakan data uji. Data uji diproses dari file *Excel* dan diumpkan ke model *XLNet* untuk menghasilkan *accuracy*. Hasil Pengujian Model *XLNet* dapat dilihat pada Gambar 11.

Test Accuracy of the finetuned model on test data is: 88.45 %

Gambar 11. Hasil Pengujian Model *XLNet*

Evaluasi *Confusion Matrix*

Evaluasi model menggunakan *confusion matrix* menghasilkan metrik *recall*, *precision*, *accuracy*, dan *f1-score* untuk menilai performa. Gambar 12 menunjukkan jumlah data berlabel positif 603, negatif 255, dan netral 112.



Gambar 12. Hasil Evaluasi

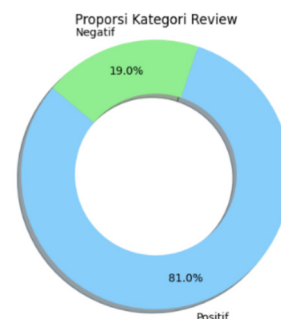
Gambar 13 menunjukkan hasil evaluasi metode *XLNet* dengan akurasi 88,45%, *precision* 89%, *recall* 88%, dan *f1-score* 89%.

	precision	recall	f1-score	support
Positif	0.93	0.90	0.92	670
Negatif	0.79	0.85	0.82	300
accuracy			0.88	970
macro avg	0.86	0.88	0.87	970
weighted avg	0.89	0.88	0.89	970

Gambar 13. *Classification Report*

Interpretasi Hasil Klasifikasi

Hasil klasifikasi *XLNet* menunjukkan 81,0% sentimen positif dan 19,0% negatif, mengindikasikan mayoritas respons positif terhadap penggunaan ChatGPT dalam pendidikan. Hasil ada pada Gambar 14.



Gambar 14. Hasil Prediksi Sentimen

Klasifikasi dengan *fine-tuning XLNet* menunjukkan akurasi 88,45%, *precision* positif 95%, *recall* 90%, dan *F1-score* 92%. *Precision* negatif 79%, *recall* 88%, dan *F1-*

- Li, H., Zhang, X., Liu, Y., Zhang, Y., Wang, Q., Zhou, X., Liu, J., Wu, H., & Wang, H. (2019). D-net: A simple framework for improving the generalization of machine reading comprehension. *MRQA@EMNLP 2019 - Proceedings of the 2nd Workshop on Machine Reading for Question Answering*, 212–219.
- Murcahyanto, H. (2023). Penerapan Media Chat GPT pada Pembelajaran Manajemen Pendidikan terhadap Kemandirian Mahasiswa. *Edumatic: Jurnal Pendidikan Informatika*, 7(1), 115–122.
<https://doi.org/10.29408/edumatic.v7i1.14073>
- Musfiroh, D., Khaira, U., Utomo, P. E. P., & Suratno, T. (2021). Analisis Sentimen terhadap Perkuliahan Daring di Indonesia dari Twitter Dataset Menggunakan InSet Lexicon. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 1(1), 24–33.
<https://doi.org/10.57152/malcom.v1i1.20>
- Pasek, P., Mahawardana, O., Sasmita, G. A., Agus, P., & Pratama, E. (2022). Analisis Sentimen Berdasarkan Opini dari Media Sosial Twitter terhadap “Figure Pemimpin” Menggunakan Python. In *JITTER-Jurnal Ilmiah Teknologi dan Komputer* (Vol. 3, Issue 1).
- Ratniasih, N. L., Jayanti, N. W. N., & ... (2023). Klasifikasi Kepuasan Mahasiswa Menggunakan Algoritma Support Vector Machine dan Metode Stemming Sastrawi. *Prosiding CORISINDO*, 373–378.
- Santoso, D. P., & Wibowo, W. (2022). Analisis Sentimen Ulasan Aplikasi Buzzbreak Menggunakan Metode Naïve Bayes Classifier pada Situs Google Play Store. *Jurnal Sains Dan Seni ITS*, 11(2).
<https://doi.org/10.12962/j23373520.v11i2.72534>
- Singh, B., Olds, T., Brinsley, J., Dumuid, D., Virgara, R., Matricciani, L., Watson, A., Szeto, K., Eglitis, E., Miatke, A., Simpson, C. E. M., Vandelanotte, C., & Maher, C. (2023). Systematic review and meta-analysis of the effectiveness of chatbots on lifestyle behaviours. *Npj Digital Medicine*, 6(1), 1–10.
<https://doi.org/10.1038/s41746-023-00856-1>
- Taecharungroj, V. (2023). “What Can ChatGPT Do?” Analyzing Early Reactions to the Innovative AI Chatbot on Twitter. *Big Data and Cognitive Computing*, 7(1).
<https://doi.org/10.3390/bdcc7010035>
- Topal, M. O., Bas, A., & van Heerden, I. (2021). *Exploring Transformers in Natural Language Generation: GPT, BERT, and XLNet*.
- Wandani, A. (2021). Sentimen Analisis Pengguna Twitter pada Event Flash Sale Menggunakan Algoritma K-NN, Random Forest, dan Naive Bayes. *Jurnal Sains Komputer & Informatika (J-SAKTI)*, 5(2), 651–665.