

ANALISIS PERBANDINGAN CONVOLUTIONAL NEURAL NETWORK DAN SUPPORT VECTOR MACHINE PADA ANALISIS SENTIMEN RUU PILKADA DI APLIKASI X

Comparative Analysis of Convolutional Neural Network and Support Vector Machine in Sentiment Analysis of the Regional Election Bill on the X Application

Yohana Uba Ladopurab¹, Yampi R Kaesmetan²

Program Studi Teknik Informatika Strata Satu¹,
Program Studi Teknik Informatika Strata Satu²
STIKOM Uyelindo Kupang¹²

yahanalpurab@gmail.com¹, kaesmetanyampi@gmail.com²

Received: 2025-03-25. **Revised:** 2025-04-28. **Accepted:** 2025-05-07. **Issue Period:** Vol.9 No.2 (2025), Pp. 908-927

Abstrak: Regulasi terkait Pemilihan Kepala Daerah (RUU Pilkada) sering menjadi perbincangan publik dan memunculkan beragam opini di media sosial, termasuk di Aplikasi X. Analisis sentimen terhadap opini publik dapat memberikan wawasan berharga bagi pembuat kebijakan dan masyarakat dalam memahami persepsi terhadap RUU Pilkada. Dalam penelitian ini, dilakukan perbandingan antara metode *Convolutional Neural Network* (CNN) dan *Support Vector Machine* (SVM) untuk analisis sentimen terhadap data yang dikumpulkan dari Aplikasi X. Berdasarkan hasil evaluasi, SVM menunjukkan kinerja yang lebih dibandingkan CNN pada hampir semua matriks, dengan akurasi SVM mencapai 92,38% dan F1-score 27,27%. Sebaliknya, CNN mencatatkan akurasi 62,86 dan F1-score 7,14%. Meskipun CNN unggul dalam pemrosesan data sekuensial, SVM terbukti lebih stabil dan akurat dalam konteks dataset ini.

Kata kunci: Analisis Sentimen, CNN, RUU Pilkada, SVM.

Abstract: Regulations related to Regional Head Elections (Regional Election Bill) Frequently become a topic of public discussion, generating diverse opinions on socia media, including the X Application. Sentiment analysis of pubic opinion can provide valuabe insights for policymakers and society in understandring perceptions of the Regional Election Bill. This study compares the *Convolutional Neural Network* (CNN) and *Support Vector Machine* (SVM) methods for sentiment analysis on data collected from the X Application. Based on the evaluation resultd, SVM demonstrated superior performance compared to CNN across nearly all metrics, achieving an accuracy of 92.38% and an F1-score of 27.27%. In contrast, CNN recorded an accuracy of 62.86% and an F1-score f 7.14%. Although CNN excels in processing sequential data, SVM proved to be more stable and accurate within the context of this dataset.

DOI: 10.52362/jisamar.v9i2.1887

Cipta



skan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

Keywords: CNN, Regional Election Bill, Sentiment Analysis, SVM

I. PENDAHULUAN

Pemilihan Kepala Daerah (Pilkada) merupakan salah satu proses demokrasi yang penting dalam sistem pemerintahan di Indonesia [1]. Sejalan dengan perkembangan ketatanegaraan dan perkembangan masyarakat pengaturan tentang Pemilihan Kepala Daerah dan Wakil Kepala Daerah, baru-baru ini Mahkamah Konstitusi (MK) memberikan putusan penting yang memengaruhi dinamika politik dan sistem pemilihan kepala daerah (Pilkada) di Indonesia. Pada 20 Agustus 2024 MK mengeluarkan keputusan terkait batas usia calon kepala daerah tertulis dalam putusan Nomor 70/PUU-XXII/2024 yang menyatakan “berusia paling rendah 30 tahun untuk Calon Gubernur dan Wakil Gubernur dan 25 tahun untuk Calon Bupati dan Wakil Bupati atau Calon Wali Kota dan Wakil Wali Kota terhitung sejak penetapan Pasangan Calon” [2].

Aplikasi X merupakan salah satu platform media sosial yang saat ini menjadi sangat populer di kalangan masyarakat Indonesia. Aplikasi ini memungkinkan penggunanya untuk berbagi berbagai informasi, termasuk pandangan, opini, dan berita terkait isu-isu terkini [3], termasuk isu-isu politik seperti rancangan undang-undang (RUU). Komentar-komentar pada twitter inilah yang sebenarnya dapat digunakan pemerintah sebagai acuan evaluasi kebijakan [4]. Popularitas dan keterbukaan platform ini menjadikannya sumber data yang kaya untuk memahami sentimen masyarakat terkait berbagai permasalahan sosial dan politik yang berkembang. Banyak riset yang memakai media twitter atau aplikasi X untuk mencari informasi salah satunya dengan analisis sentimen yang digunakan untuk menganalisis opini yang ada dalam suatu topik di twitter [5].

Penelitian ini, akan menganalisa sentimen publik pada platform Twitter atau aplikasi X menjadi relevan untuk memahami reaksi, pandangan, dan perasaan masyarakat terhadap masalah yang terjadi [6] Analisis sentimen merupakan jenis pemrosesan bahasa alami yang digunakan untuk mengidentifikasi suasana hati seseorang terkait topik tertentu [7]. Analisis sentimen akan memfasilitasi ekstraksi beragam pendapat dari Twitter atau aplikasi X. Temuan ini akan mengkonfirmasi apakah sentimen yang diungkapkan itu baik, negatif, atau netral[8]. Keberadaan analisis sentimen, mengatasi persoalan-persoalan dengan cara mengubah format yang tidak terstruktur menjadi format terstruktur dan terklasifikasi [9].

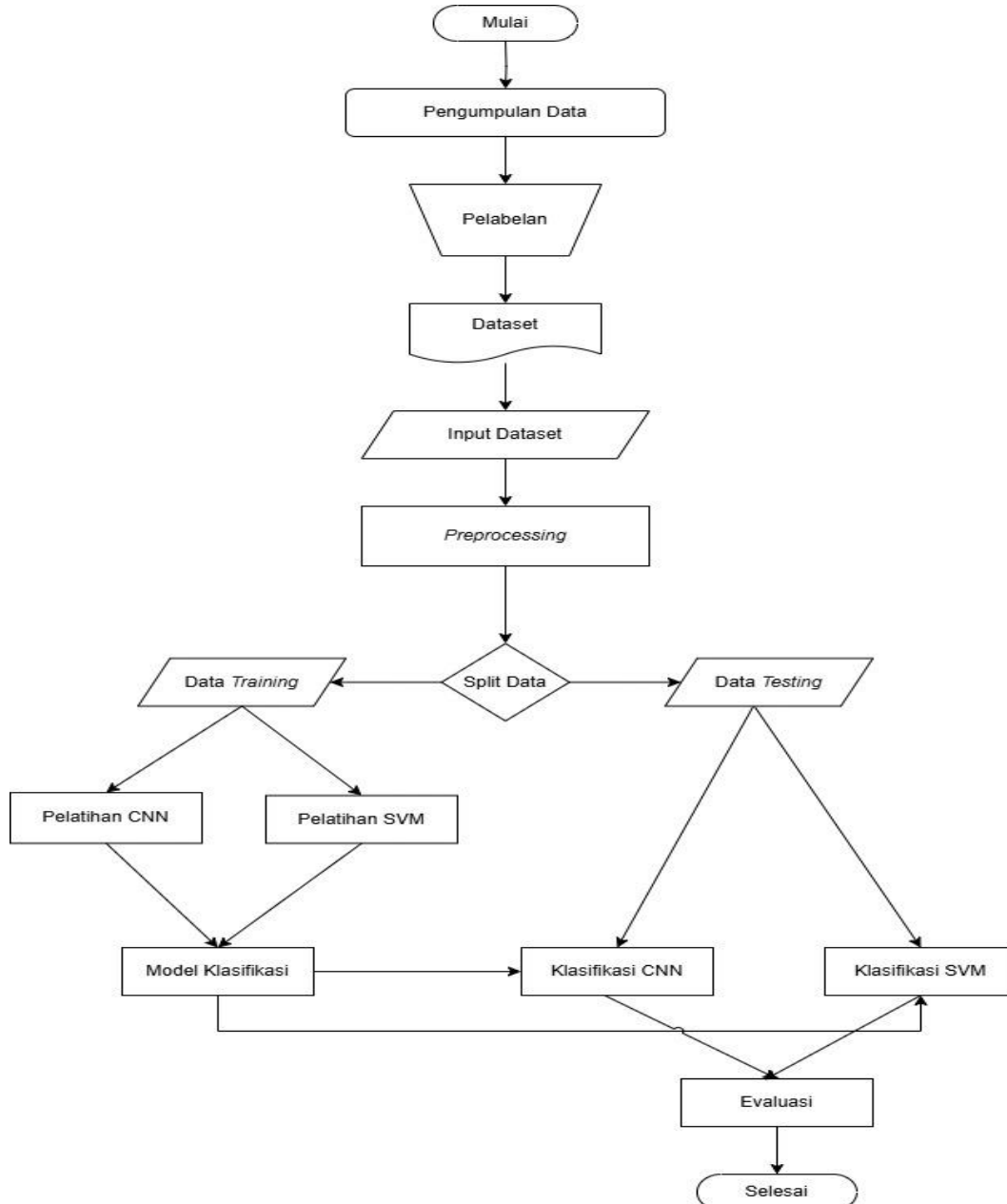
Berdasarkan pembahasan di atas, penelitian ini diharapkan dapat memberikan kontribusi dalam memahami persepsi publik terhadap RUU Pilkada melalui analisis data media sosial. Penelitian ini akan menggunakan metode Convolutional Neural Network (CNN) dan Support Vector Machine (SVM) untuk menganalisis sentimen yang ada. Analisis ini juga diharapkan dapat menjadi referensi bagi penelitian selanjutnya yang berkaitan dengan penerapan CNN dan SVM dalam analisis sentimen di media sosial.



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

II. LITERATUR DAN METODE



Gambar 1. Prosedur Penelitian

Berdasarkan prosedur penelitian, perancangan alur setiap proses dalam sistem yang akan dibangun dapat dijelaskan sebagai berikut. Tahap awal dimulai dengan pengambilan data berupa respons masyarakat Indonesia terhadap RUU Pilkada melalui Aplikasi X, yang dilakukan menggunakan platform *Google Colab*. Data yang telah diperoleh kemudian diberi label sesuai dengan sentimen masing-masing ulasan, apakah termasuk dalam kategori positif dan negatif. Setelah proses pelabelan, data tersebut selanjutnya melalui tahap prapemrosesan (*preprocessing*) yang bertujuan untuk

membersihkan dan mempersiapkan data agar layak digunakan dalam proses pemodelan. Setelah data selesai diproses, dilakukan pembagian data secara acak menjadi dua bagian, yaitu 80% sebagai data latih (*training*) dan 20% sebagai data uji (*testing*). Tahapan berikutnya adalah pemberian bobot pada setiap term yang terdapat dalam data latih menggunakan metode *Term Frequency-Inverse Document Frequency (TF-IDF)*, guna merepresentasikan tingkat kepentingan masing-masing term dalam dokumen terhadap keseluruhan korpus. Setelah proses pembobotan *term*, data *training* kemudian akan dilatih menggunakan metode CNN dan SVM. Dari hasil penerapan dari kedua metode tersebut kemudian dilakukan pengujian dengan *confussion matrix* yang kemudian menghasilkan nilai *accuracy*, *precision*, *recall* dan *specifity*.

A. Media Sosial X

Media sosial adalah sebuah media online. Para penggunanya bisa dengan mudah berpartisipasi, berbagi, dan menciptakan isi meliputi blog, jejaring sosial, wiki, forum, dan dunia virtual [10]. Aplikasi X merupakan platform media sosial di mana orang dapat berbagi pesan singkat, gambar, dan video. Platform ini telah berkembang menjadi tempat di mana orang dari seluruh dunia berkumpul untuk membahas ide-ide baru dan tren [11]. Aplikasi X menjadi urutan ke 5 sebagai media sosial terfavorit di Indonesia pada tahun 2022 dengan jumlah penggunanya sejumlah 14,85 juta [12].

B. RUU Pilkada

RUU Pilkada merupakan salah satu rancangan reguasi yang bertujuan untuk mengatur mekanisme dan prosedur pelaksanaan Pemilihan Kepala Daerah di Indonesia. Pilkada memiliki peran penting dalam mendemokratisasikan kehidupan politik daerah karena menentukan siapa yang akan memimpin pemerintahan. Namun, proses tersebut tidak terlepas dari berbagai permasalahan hukum dan prosedural yang terkadang menjadi bahan perdebatan dan kontroversi [1].

Revisi Undang-Undang yang dilakukan oleh DPR menghasilkan putusan yang berbeda dengan putusan MK padahal menurut Undang-Undang Dasar Negara Republik Indonesia 1945 Pasal 24 C, yang menyatakan bahwa putusan Mahkamah Konstitusi bersifat final untuk menguji Undang-Undang [2]. Hal ini memicu terjadinya perdebatan umum yang ditandai dengan aksi demonstrasi di setiap daerah. Bukan hanya sekedar demonstrasi namun aksi protes ini juga terjadi di berbagai platform media sosial seperti pada Aplikasi X.

C. Klasifikasi

Klasifikasi merupakan teknik untuk mengklasifikasikan suatu data tertentu ke dalam kelas yang sebelumnya telah ditentukan kelasnya pada pola atau atribut yang terdapat pada data tersebut [13]. Klasifikasi melakukan pembangunan model berdasarkan data latih yang ada, kemudian menggunakan model tersebut untuk mengklasifikasikan pada data yang baru [14]. Tujuan dari klasifikasi merupakan untuk memprediksi ataupun memperkirakan kelas dari satu data terkini yang belum mempunyai merek [15]. Klasifikasi merupakan proses yang terdiri dari dua tahap, yaitu tahap pembeajaran dan tahap pengklasifikasian. Pada tahap pembelajaran, sebuah algoritma klasifikasi akan membangun sebuah model klasifikasi dengan cara menganalisis training data. Tahap pembelajaran dapat juga dipandang sebagai tahap pembentukan fungsi atau pemetaan $Y=F(X)$ dimana Y adalah kelas hasil prediksi dan X adalah *tuple* yang ingin di diprediksi kelasnya. Selanjutnya pada tahap pengklasifikasian, model yang dihasilkan akan digunakan untuk melakukan pengklasifikasian [16].

D. Opinion Mining

Opinion mining merupakan teknik yang digunakan untuk menganalisis dan mengklasifikasikan ulasan dari kata, kalimat, maupun dokumen dengan tujuan untuk memahami sentimen atau pendapat yang terkandung di dalamnya [17]. *Opinion mining* dapat dimanfaatkan untuk mengetahui respon pengguna media sosial terhadap suatu isu, sehingga apabila diperlukan pengambilan kebijakan bagi banyak pihak, hasil dari *opinion mining* dapat dijadikan bahan pertimbangan [18].

E. Text Mining



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

Text mining secara luas dapat didefinisikan sebagai intensif dimana *user* berinteraksi dengan suatu kumpulan dokumen dari waktu ke waktu dengan menggunakan seperangkat alat analisis. Penambahan teks berupaya mengekstraksi informasi yang berguna dari sumber data melalui identifikasi dan eksplorasi dengan pola yang menarik [19].

F. Pre-Processing Data

Pre-processing bertujuan menghilangkan kata-kata yang tidak bermakna. Proses pembersihan membuat data lebih teragregasi, memudahkan penggunaan dalam tahap berikutnya [20]. Pada pemrosesan selanjutnya, *preprocessing* dataset melibatkan beberapa langkah seperti:

1. Seleksi Data
Dalam tahap ini dilakukan pemilihan kolom atribut dimana hanya kolom *text* saja sebagai opini masyarakat yang akan dikelola pada tahap berikutnya.
2. *Cleaning Data*
Cleaning data membersihkan data "*full_text*" yang berisi *tweet* dari *url*, *html*, *emoji*, penomoran, dan simbol [17].
3. *Case Folding*
Case Folding merupakan metode untuk menyederhanakan data teks, menormalisasi teks dengan mengubah seluruh teks yang semula menjadi huruf acak menjadi huruf kecil hanya huruf A-Z [21].
4. Normalisasi
Normalisasi adalah proses penskalaan nilai atribut dari data sehingga bisa terletak pada rentang tertentu [14].
5. *Process Document from Data*
Rangkaian beberapa operator sub proses dalam operator *process document from data* pada rangkaian utama, yang digunakan *tokenize*, *transform case*, *filter stopword*, *filter token by length* [22].
 - a. *Tokenize*
Tahap ini adalah tahap dimana pemotongan kata berdasarkan tiap kata yang penyusunannya menjadi potongan tunggal. Kata dalam kata yang dimaksud adalah kata yang dipisah oleh spasi.
 - b. *Transform Cases*
Tahap ini berfungsi untuk menyamakan semua huruf, menjadi huruf kecil semua.
 - c. *Filter Stopword*
Tahap ini merupakan penghapusan kata yang tidak memiliki arti berdasarkan *stopwords* atau kata yang tidak memiliki makna sendiri atau tidak terkait dengan kata sifat yang berhubungan dengan *stopword*.
 - d. *Filter Token (by Length)*
Tahap ini merupakan penghapusan kata yang tidak memiliki arti berdasarkan panjang karakter yang ditentukan.

G. Pembobotan Term

Pembobotan term merupakan proses perhitungan bobot tiap term yang dicari pada setiap dokumen sehingga dapat diketahui ketersediaan dan kemiripan suatu term di dalam dokumen [23]. Salah satu teknik pembobotan adalah dengan *Term Frequency-Inverse Document Frequency* (TF-IDF) [24]. *Term Frequency-Inverse Document Frequency* adalah metode algoritma yang berguna untuk menghitung bobot setiap kata yang sering digunakan. Metode ini dikenal efisien, mudah diterapkan, dan memberikan hasil yang akurat [25].

H. Convolutional Neural Network

Convolutional Neural Network merupakan bagian dari salah satu jenis arsitektur *deep learning* yang sering digunakan dalam pengolahan data gambar dan teks. CNN termasuk komponen dari jaringan saraf mendalam (*Deep Neural Network*) [21] CNN memiliki proses dari awal masukan melewati *convolution layer*, *pooling layer*, dan *fully connected layer* [26].

I. Support Vector Machine

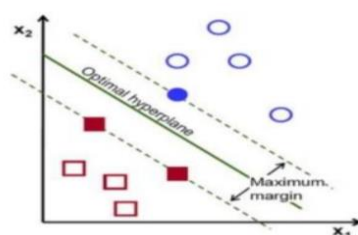
DOI: 10.52362/jisamar.v9i2.1887

Cipta



skan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

Support Vector Machine (SVM) adalah algoritma klasifikasi yang pertama kali diutarakan oleh Boser, Guyon serta Vapnik pada 1992. SVM adalah metode klasifikasi yang menggunakan metode *machine learning* (*supervised learning*) dan meramalkan kelas berdasarkan pola [27]. SVM memiliki cara khas berusaha menyediakan jarak kelas maksimum ke data, istilahkan dengan *hyperplane*, seperti yang derlihatkan pada gambar 1. Trik kernel diterapkan untuk meningkatkan jarak antar kelas tanpa memberikan beban komputasi yang tinggi [28]. [29].



Gambar 2. Ilustrasi SVM

J. Confussion Matrix

Confussion Matrix adalah metode yang digunakan untuk menghitung akurasi dalam *Data Mining*. Evaluasi menggunakan *Confussion Matrix* menghasilkan nilai akurasi, presisi, *recall* [30].

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Gambar 3. *Confussion Matrix*

K. Python

Python adalah bahasa pemrograman interpretatif yang mudah dipelajari dan dapat dijalankan pada berbagai platform dengan fokus utama pada keterbacaan kode [31]. *Python* adalah bahasa pemrograman yang memanfaatkan interpreter untuk menjalankan kode program secara langsung. Interpreter ini mampu menerjemahkan kode tanpa perlu mengompilasinya terlebih dahulu [29]. Banyaknya tweet yang masuk terkait RUU Pilkada mendorong perlunya metode baru untuk melihat opini masyarakat secara efektif. *Python* merupakan bahasa pemrograman yang dapat digunakan untuk menjawab permasalahan tersebut [32].

III. HASIL DAN PEMBAHASAN

A. Crawling Data

Crawling data yang diambil dari media sosial X dengan kata kunci “RUU Pilkada” adalah proses pengumpulan informasi yang dilakukan secara otomatis untuk mendapatkan data dalam bentuk teks. Dalam hal ini, teks yang dikumpulkan adalah *tweet* dari pengguna Aplikasi X. Pengumpulan data ini dilakukan dalam rentang waktu yang cukup lama, yakni mulai dari tanggal 28 Agustus 2024 hingga 2 Maret 2025. Selama periode tersebut, ada 1061 *tweet* yang berhasil diambil.

Namun, tidak semua *tweet* tersebut dapat langsung digunakan untuk analisis lebih lanjut. Oleh karena itu, dilakukan tahap *pre-processing* yang meliputi berbagai proses seperti penghapusan karakter khusus, normalisasi teks dan pemilihan *tweet* yang relevan dengan topik yang sedang diteliti. Setelah melalui tahap *cleaning*, jumlah *tweet* yang siap



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

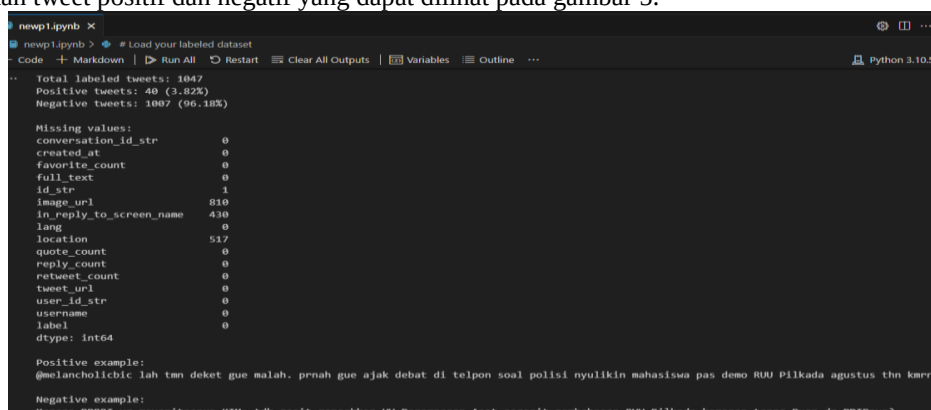
untuk dianalisis berkurang menjadi 1047 *tweet*. *Tweet-tweet* ini kemudian diklasifikasikan berdasarkan sentimen menjadi dua kategori, yaitu opini positif dan opini negatif. Hasil *crawling* data dapat dilihat pada gambar 4.

	1	conversat	created_a	favorite_c	full_text	id_str	image_url	in_reply_t	lang	location	quote_co	reply_cou	retweet_c	tweet_url	user_id_s	username
2	1.90E+18	Sun Mar 0	0	Kenapa Di	1.90E+18			in		Sangatta	0	0	0	https://x.	1.96E+09	ferdybakti
3	1.90E+18	Fri Feb 28	0	@mkusun	1.90E+18			in			0	0	0	https://x.	1.19E+18	afdiyasrul
4	1.89E+18	Thu Feb 2	0	@ahmads	1.90E+18			in			0	1	0	https://x.	1.81E+18	m_harry59032
5	1.89E+18	Thu Feb 2	4	dulu pas c	1.89E+18			in		neverland	0	1	1	https://x.	1.05E+18	jijirnyinyir
6	1.89E+18	Thu Feb 2	2	@gearcloi	1.89E+18			in		neverland	0	0	0	https://x.	1.05E+18	jijirnyinyir
7	1.89E+18	Thu Feb 2	0	Dan PDI-P	1.89E+18			in		RPLLegaw	0	0	0	https://x.	1.75E+18	RPLLegaw
8	1.89E+18	Thu Feb 2	13	Terima ka	1.89E+18			in		Jakarta	0	2	3	https://x.	59387308	titiangraini
9	1.89E+18	Wed Feb 2	0	@akbarfai	1.89E+18			in		BSD - Mat	0	0	0	https://x.	1.85E+09	riddhopratama
10	1.89E+18	Wed Feb 2	0	Badan Legislati	1.89E+18			in			0	0	0	https://x.	1.83E+18	SuaraTrenggalek
11	1.89E+18	Mon Feb 2	0	Terus ujug	1.89E+18			in		Jakarta Ba	0	1	0	https://x.	3.27E+09	matilusa
12	1.89E+18	Mon Feb 2	0	@EDH012	1.89E+18			in			0	1	0	https://x.	1.18E+18	tekarok007
13	1.89E+18	Sun Feb 2	0	@MbulGe	1.89E+18			in			0	2	0	https://x.	4.53E+08	azzamrabbani
14	1.89E+18	Sun Feb 2	5	@GOLIAH8	1.89E+18			in			0	0	0	https://x.	1.28E+18	judee_hey
15	1.89E+18	Sun Feb 2	0	Ahmad Dc	1.89E+18			in			0	0	0	https://x.	8.54E+17	radaraktual
16	1.89E+18	Sat Feb 22	16	Perbandir	1.89E+18		https://pbs.twimg.c	in			0	2	4	https://x.	1.76E+18	platonoyunsky
17	1.89E+18	Sat Feb 22	3	@Justforf	1.89E+18			in		DKI Jakart	0	1	0	https://x.	7.72E+08	CarolineSarah5
18	1.89E+18	Sat Feb 22	0	@PartaiSc	1.89E+18			in			0	0	0	https://x.	1.27E+18	dhimasairlangg1
19	1.89E+18	Sat Feb 22	5	Ahmad Dc	1.89E+18			in		DKI Jakart	0	0	35	https://x.	7.22E+17	golkarpedia
20	1.89E+18	Fri Feb 21	6	@AndrePi	1.89E+18			in			0	0	0	https://x.	1.75E+18	Roselly28
21	1.89E+18	Fri Feb 21	0	Baleg DPR	1.89E+18			in			0	0	0	https://x.	7.29E+17	TajukNasional

Gambar 4. Hasil *Crawling*

B. Pelabelan Dataset

Pada proses ini dilakukan persiapan dataset dengan membaca data yang kemudian dilakukan distribusi sentimen untuk mengetahui jumlah *tweet* positif dan negatif yang dapat dilihat pada gambar 5.



```

new1.ipynb
# Load your labeled dataset
Code + Markdown ▶ Run All ⌂ Restart Clear All Outputs Variables Outline ... Python 3.10.5

Total labeled tweets: 1047
Positive tweets: 40 (3.82%)
Negative tweets: 1007 (96.18%)

Missing values:
conversation_id_str    0
created_at             0
favorite_count         0
full_text              0
id_str                 1
image_url              810
in_reply_to_screen_name 430
lang                   0
location               517
quote_count            0
reply_count            0
retweet_count          0
tweet_url              0
user_id_str            0
username               0
label                  0
dtype: int64

Positive example:
@melancholicbic lah tmn dekat gue malah. prnah gue ajak debat di telpon soal polisi nyulikin mahasiswa pas demo RUU Pilkada agustus tmn kmr

Negative example:
Kenapa DPRRI ya mayoritasnya KTM? tdk gesit mensahkan RUU Perampasan Aset segesit pembahasan RUU Pilkada kemarin tanpa Puan dan PDIPnya?

```

Gambar 5. Hasil Pelabelan

Pada gambar di atas menunjukkan hasil dari proses pelabelan dimana dari 1047 *tweet* diperoleh 40 *tweet* positif dan 1007 *tweet* negatif. Data tidak memiliki nilai hilang pada kolom-kolom penting seperti *full-text* dan *label*, sehingga dapat langsung digunakan untuk analisis sentimen. Beberapa kolom seperti *image_url*, *location* dan *in_reply_to_screen_name* memang memiliki banyak nilai kosong, yang merupakan hal umum dalam data X.

C. Precession

Pada *preprocessing* ini dilakukan *cleaning* dan *case folding* untuk menghapus *noise* seperti *url*, *mention*, *hashtag*, angka dan kemudian menyamakan format. Kemudian dilanjutkan dengan menyederhanakan kata dan menghilangkan kata umum yang biasa disebut normalisasi dan *stopword*. Setelah itu dilakukan proses *stemming* untuk mengembalikan ke akar kata untuk generalisasi. Berikut gambar 6 adalah hasil pembersihan data dan pra-pemrosesan pada dataset.

```
[nltk_data] Downloading package punkt to
[nltk_data]   C:\Users\ACER\AppData\Roaming\nltk_data...
[nltk_data]   Package punkt is already up-to-date!
[nltk_data] Downloading package stopwords to
[nltk_data]   C:\Users\ACER\AppData\Roaming\nltk_data...
[nltk_data]   Package stopwords is already up-to-date!

Starting text preprocessing...

Preprocessed examples:
Original: Kenapa DPRRI yg mayoritasnya KIM* tdk gesit mensahkan UU Perampasan Aset segesit pembahasan RUU Pilik...
Processed: dprri mayoritas kim gesit sah uu ampas aset gesit bahas ruu pilkada kemaren puan dg pdipnya...
Label: 0

Original: @mkusumawijaya tp ruu pilkada bisa dibidang berhasil walaupun ttp 50:50...
Processed: tp ruu pilkada bidang hasil ttp...
Label: 0

Original: @ahmadsyahrul_03 @tanyakanrl Kalau demo cuma ampe sore itu PPN tetep 12% tolol. Itu RUU Pilkada temp...
Processed: demo ampe sore ppn tetep tolol ruu pilkada tempo batal tolol...
Label: 0

Preprocessing completed. Samples saved to output/data/preprocessing_samples.csv
```

Gambar 6. Hasil Pembersihan Data dan Pra-Pemrosesan Teks

Berdasarkan gambar 6, setiap *tweet* telah dibersihkan dari kata-kata tidak penting, simbol, dan kemungkinan dilakukan normalisasi kata. Hasil akhir terlihat lebih ringkas dan berfokus pada kata-kata kunci yang relevan untuk analisis sentimen.

D. Feature Extraction Using TF-IDF

Pada proses ini dilakukan untuk menghitung banyaknya *term* atau kata yang muncul pada *tweet* (tf0, menghitung banyaknya *tweet* yang mengandung *term* tersebut (df), menghitung inverse dokumen *frequency* (idf) dan mengalikan tf dengan idf sebagai bobot dari *term* pada setia *tweet*. Hasil dari proses ekstraksi fitur menggunakan TF-IDF diperlihatkan oleh gambar 7.

```
Starting feature extraction with TF-IDF...
Data split statistics:
  Dataset  Size  Positive Samples  Negative Samples
0 Training Set  837      32      805
1 Testing Set   210       8      202

Top 20 Features by TF-IDF sum:
ruu: 47.5016
pilkada: 45.1109
mk: 41.5767
putus: 41.5418
kawal: 36.7042
demo: 29.0091
aksi: 24.4325
dpr: 20.3478
aman: 16.2200
apresiasi: 15.7602
polri: 15.5666
bobby: 15.4645
semmi: 15.1873
pb: 15.1873
kurniawan: 15.1873
mp: 15.1873
rakyat: 14.0098
kemarin: 13.5340
aja: 13.5034
dukung: 12.3053

Feature extraction completed and saved to output directory.
```

Gambar 7. Hasil Feature Extraction Using TF-IDF

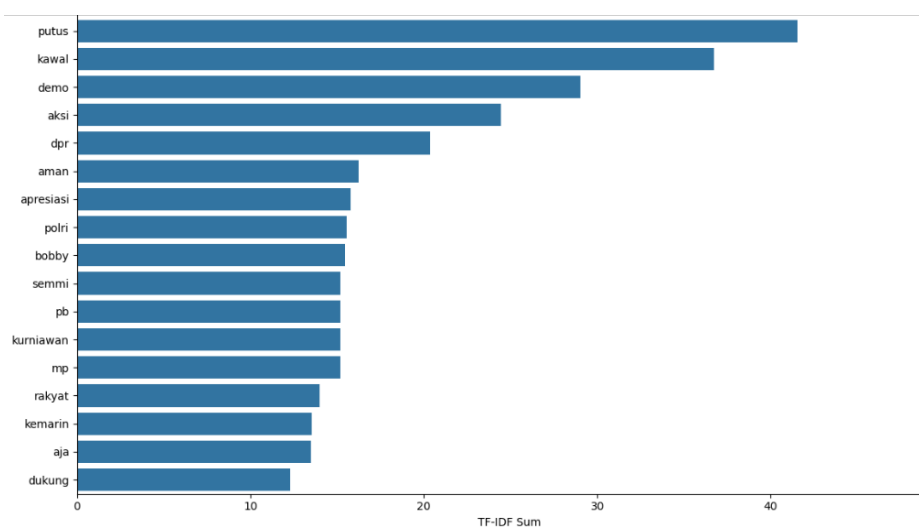
Gambar 7 merupakan output dari proses ekstraksi fitur menggunakan TF-IDF dimana data telah dibagi menjadi dua set yaitu 837 data latih dan 210 data uji. Dari hasil perhitungan TF-IDF, kata-kata yang paling menonjol dalam korpus *tweet* antara lain adalah “ruu”, “pilkada”, “mk”, “putus”, “kawal” dan yang lainnya menunjukkan fokus topik pada isu-isu legislatif dan aksi masyarakat.

Dengan merujuk pada 20 kata dengan skor TF-IDF tertinggi di dataset yang memberikan wawasan tentang kata-kata yang paling berpengaruh. Kata yang dimaksud ditunjukkan pada gambar 8.



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).



Gambar 8. Frekuensi Teks

E. Analisa Sentimen dengan Algoritma SVM

Data yang sudah melalui *preprocessing* kemudian menggunakan SVM untuk menentukan parameter terbaik dan dievaluasi menggunakan *confusion matrix*. Hasil pemrosesan menggunakan svm dan hasil evaluasi dapat dilihat ada gambar 9 dan 10.

```
Starting SVM model training and evaluation...
Training SVM model...
Fitting 5 folds for each of 3 candidates, totalling 15 fits
Best parameters: {'C': 0.1, 'kernel': 'linear'}
SVM Model Evaluation:
Accuracy: 0.9238
Precision: 0.2143
Recall: 0.3750
Specificity: 0.9455
F1 Score: 0.2727

Classification Report:
      precision    recall  f1-score   support

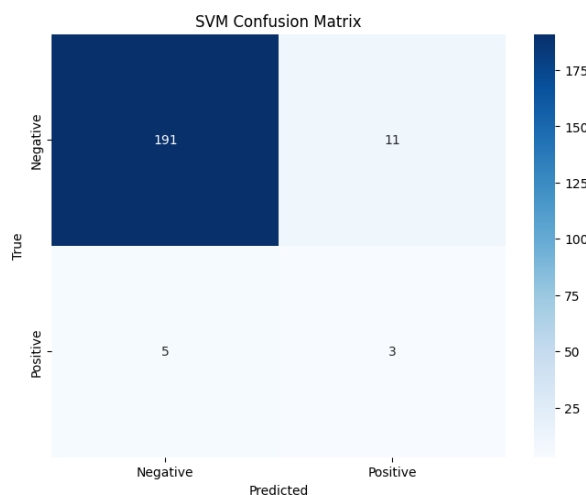
     0       0.97       0.95       0.96        202
     1       0.21       0.38       0.27         8

   accuracy       0.92        210
  macro avg       0.59       0.66       0.62        210
 weighted avg       0.95       0.92       0.93        210

Confusion Matrix:
[[191  11]
 [  5   3]]
```

Gambar 9. Model SVM

Pada gambar 9 terlihat bahwa model *Support Vector Machine (SVM)* berhasil dilatih dengan parameter terbaik yaitu kernel linear dan nilai C sebesar 0.1. Hasil menunjukkan akurasi tinggi sebesar 92,28%, namun performa dalam mendeteksi *tweet* positif masih rendah, dengan *precision* sebesar 21,43%, *recall* 37,5% dan *f1-score* 27,27%.



Gambar 10. Confusion Matrix SVM

Gambar di atas menunjukkan *confusion matrix* dari model SVM. Baris pertama mempresentasikan *tweet* dengan label sebenarnya negatif, dimana 191 *tweet* berhasil diprediksi dengan benar sebagai negatif, dan 11 *tweet* diprediksi sebagai positif. Baris kedua menunjukkan *tweet* dengan label sebenarnya positif, dimana 3 *tweet* diprediksi dengan benar sebagai benar dan 5 *tweet* lainnya diprediksi sebagai negatif.

F. Analisa Sentimen dengan Algoritma CNN

Data yang sudah melalui *preprocessing* kemudian dilatih dengan pembobotan kelas yang seimbang untuk menangani data tidak seimbang, serta dievaluasi secara menyeluruh menggunakan *confusion matrix*. Berikut hasil pemrosesan data dan evaluasi dapat dilihat pada gambar-gambar berikut.

```
Starting CNN model training and evaluation
Tokenizing and padding sequences...
Sequence length statistics: Max=47, Mean=17.6, Using=32
X_train_pad shape: (837, 32)
Label distribution in training data: [805 32]

Distribusi label asli:
label
0    0.961768
1    0.038232
Name: proportion, dtype: float64

Distribusi label CNN (setelah konversi):
0    0.961768
1    0.038232
Name: proportion, dtype: float64
Class weights yang direvisi: {0: np.float64(0.5198757763975155), 1: np.float64(13.078125)}
Building high-performance CNN model...
c:\Users\ACER\AppData\Local\Programs\Python\Python310\lib\site-packages\keras\src\layers\core\embedding.py:90: UserWarning: Argument `input
warnings.warn(

Model Summary:
Model: "functional"
```

Gambar 11. Hasil Label Distribusi

Pada gambar 11 proses pelatihan model CNN dimulai dengan tokenisasi dan padding data teks menjadi panjang urutan 32, dengan statistik menunjukkan panjang urutan maksimum 47 dan rata-rata 17,6. Distribusi label data pelatihan sangat tidak seimbang, dengan 96% data berlabel 0 dan 3,8% berlabel 1. Untuk mengatasi ketidakseimbangan ini, digunakan *class weights*, dimana kelas 1 diberikan bobot lebih tinggi (13.08) agar model lebih memperhatikan kelas tersebut.



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

Layer (type)	Output Shape	Param #	Connected to
input_layer (InputLayer)	(None, 32)	0	-
embedding (Embedding)	(None, 32, 200)	679,000	input_layer[0][0]
spatial_dropout1d (SpatialDropout1D)	(None, 32, 200)	0	embedding[0][0]
conv1d (Conv1D)	(None, 32, 128)	51,328	spatial_dropout1...
conv1d_1 (Conv1D)	(None, 32, 128)	76,928	spatial_dropout1...
conv1d_2 (Conv1D)	(None, 32, 128)	102,528	spatial_dropout1...
conv1d_3 (Conv1D)	(None, 32, 128)	128,128	spatial_dropout1...
batch_normalization (BatchNormalizatio...	(None, 32, 128)	512	conv1d[0][0]
batch_normalization_1 (BatchNormalizatio...	(None, 32, 128)	512	conv1d_1[0][0]
batch_normalization_2 (BatchNormalizatio...	(None, 32, 128)	512	conv1d_2[0][0]
batch_normalization_3 (BatchNormalizatio...	(None, 32, 128)	512	conv1d_3[0][0]

global_max_pooling_1 (GlobalMaxPooling1...	(None, 128)	0	batch_normalizat...
global_max_pooling_2 (GlobalMaxPooling1...	(None, 128)	0	batch_normalizat...
global_max_pooling_3 (GlobalMaxPooling1...	(None, 128)	0	batch_normalizat...
global_max_pooling_4 (GlobalMaxPooling1...	(None, 128)	0	batch_normalizat...
concatenate (Concatenate)	(None, 512)	0	global_max_pooli...
dense (Dense)	(None, 256)	131,392	concatenate[0][0]
batch_normalization_4 (BatchNormalizatio...	(None, 256)	1,024	dense[0][0]
dropout (Dropout)	(None, 256)	0	batch_normalizat...
dense_1 (Dense)	(None, 128)	32,608	dropout[0][0]
batch_normalization_5 (BatchNormalizatio...	(None, 128)	512	dense_1[0][0]
dropout_1 (Dropout)	(None, 128)	0	batch_normalizat...
dense_2 (Dense)	(None, 1)	129	dropout_1[0][0]

Gambar 12. Gambar Model CNN

Gambar diatas adalah output dari model.summary() pada arsitektur model CNN di Tensorflow. Pada output tersebut terdapat nama layer dan jenis layer yang digunakan, bentuk output dari layer setelah menerima input, jumlah parameter yang bisa dilatih di layer tersebut dan layer sebelumnya yang menghubungkan data ke layer berikutnya.

```

Total params: 1,285,849 (4.60 MB)

Trainable params: 1,284,857 (4.59 MB)

Non-trainable params: 1,792 (7.00 KB)

```

Gambar 13. Jumlah Parameter

Gambar 13 menunjukkan jumlah parameter model dimana terdapat jumlah keseluruhan model baik yang dilatih maupun tidak, parameter yang akan diupdate selama proses training dan parameter yang tidak diupdate saat training.

```

Final CNN Model Evaluation with Optimized Threshold:
Accuracy: 0.6286 (62.86%)
Precision: 0.0395 (3.95%)
Recall: 0.3750 (37.50%)
Specificity: 0.6386 (63.86%)
F1 Score: 0.0714 (7.14%)

Classification Report:
              precision    recall  f1-score   support

      0       0.96       0.64       0.77        202
      1       0.04       0.38       0.07         8

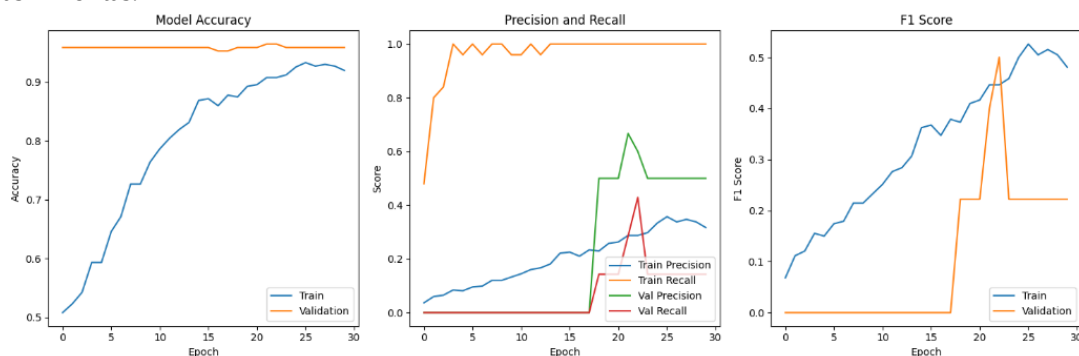
   accuracy          0.50
  macro avg          0.50
 weighted avg          0.93

Confusion Matrix:
[[129  73]
 [  5   3]]

```

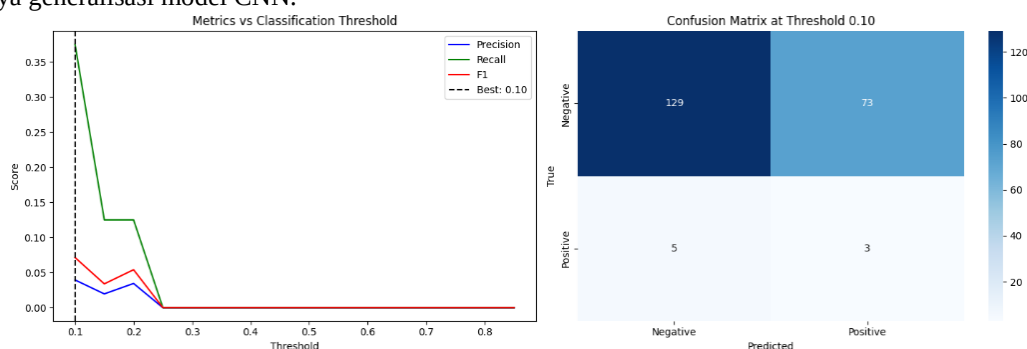
Gambar 14. Hasil Evaluasi Model CNN

Gambar 14 menunjukkan hasil evaluasi akhir model CNN dengan *threshold* yang telah dioptimasi dimana model mencapai akurasi tinggi 96.67% dan *precision* sempurna 100%, namun *recall* rendah 12.5% mengindikasikan masih banyak data positif yang tidak terdeteksi. *F1-score* juga rendah yaitu 22.22% karena ketidakseimbangan data, meskipun spesifikasi mencapai 100%. Model cenderung sangat akurat untuk kelas mayoritas, tapi kurang efektif dalam menangkap kelas minoritas.



Gambar 15. Model Accuracy, Precision, Recall dan F1 Score

Grafik pertama (*model accuracy*) menunjukkan bahwa garis biru (akurasi pelatihan) terus meningkat hingga mendekati 1, sementara garis oranye (akurasi validasi) tetap datar di sekitar 0.93, menandakan overfitting karena model hanya bagus di data pelatihan. Grafik kedua (*precision and recall*) memperlihatkan garis merah dan hijau (*precision* dan *recall* validasi) berada di angka 0 hampir sepanjang waktu, hanya naik sedikit menjelang akhir, sementara garis biru dan oranye (*precision* dan *recall* pelatihan) perlahan meningkat menunjukkan model tidak mampu mengenali kelas positif dengan baik di data validasi. Grafik ketiga (*F1-score*) menampilkan garis biru (*F1* pelatihan) yang naik stabil, sedangkan garis oranye (*F1* validasi) tetap sangat rendah, menggambarkan ketidakseimbangan performa antara pelatihan dan validasi serta rendahnya generalisasi model CNN.



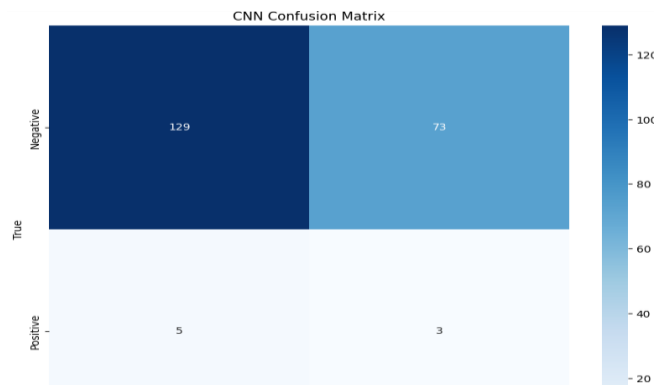
Gambar 16. Performa Model CNN

Grafik diatas menunjukkan performa model CNN terhadap berbagai nilai *threshold* klasifikasi. Sumbu X (*Threshold*) merupakan nilai ambang untuk memutuskan apakah prediksi output termasuk kelas positif atau negatif. Misalnya, jika *threshold* = 0.1 maka prediksi ≥ 0.1 dianggap positif. Sedangkan sumbu Y (*Score*) merupakan nilai metrik evaluasi (*precision*, *recall* dan *F1-score*) pada tiap *threshold*. Pada grafik juga terdapat garis berwarna yang mana garis biru melambangkan *precision*, garis hijau untuk *recall*, garis merah untuk *F1-score* dan garis putus-putus pada 0.10 yang menghasilkan *F1-score* tertinggi sebagai *threshold* terbaik. Pada *confusion matrix* pada *threshold* 0.10 dapat dilihat nilai dari TN, FP, FN dan TP yang jika dihitung menghasilkan nilai matrik evaluasi yang sama seperti pada grafik.



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).



Gambar 17. Confussion Matrix CNN

Gambar ini menunjukkan bagaimana model CNN (*Convolutional Neural Network*) melakukan prediksi terhadap data uji dalam klasifikasi sentimen.

G. Perbandingan Kinerja Model SVM dan CNN

Setelah data selesai dikelolah menggunakan model SVM dan juga CNN kemudian dilakukan perbandingan antara kedua model tersebut dengan cara yang sangat terperinci, baik dari sisi matriks evaluasi maupun analisis fitur, dan memberikan gambaran visual yang membantu untuk memilih model yang lebih baik. Perbandingan dilakukan berdasarkan matriks standar dan visualisasi yang membantu untuk menilai kinerja model secara lebih mudah. Hasil perbandingan dapat dilihat pada gambar 18.

```

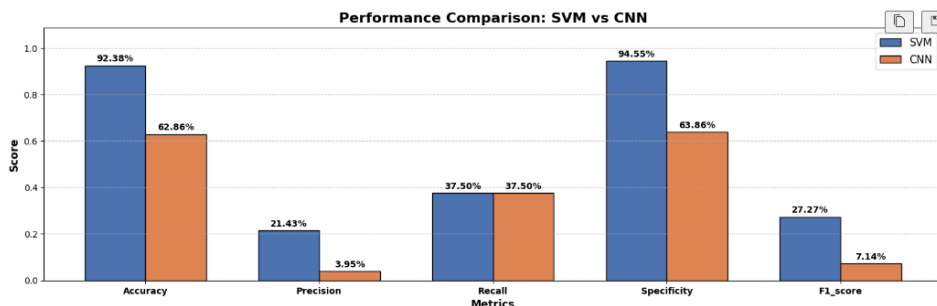
Model Performance Comparison:
Metric      SVM      CNN      Absolute Diff.  Relative Improvement
Accuracy 0.9238 (92.38%) 0.6286 (62.86%) -0.2952 -31.96%
Precision 0.2143 (21.43%) 0.0395 (3.95%) -0.1748 -81.58%
Recall 0.3750 (37.50%) 0.3750 (37.50%) 0.0000 0.00%
Specificity 0.9455 (94.55%) 0.6386 (63.86%) -0.3069 -32.46%
F1_score 0.2727 (27.27%) 0.0714 (7.14%) -0.2013 -73.81%

Overall Performance (Simple Average):
SVM: 0.5463 (54.63%)
CNN: 0.3506 (35.06%)

Overall Performance (Weighted Average):
SVM: 0.5419 (54.19%)
CNN: 0.3360 (33.60%)

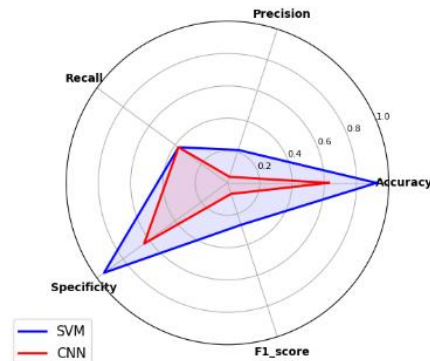
The SVM model performed better with a weighted score of 0.5419 (54.19%)
  
```

Gambar 18. Hasil Perbandingan



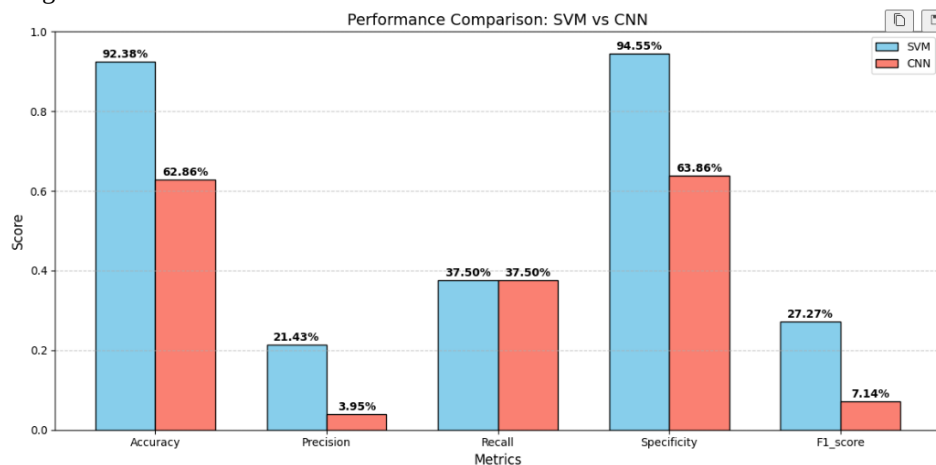
Gambar 19. Diagram Batang Hasil Perbandingan

Grafik yang ditunjukkan pada gambar 19 adalah perbandingan kinerja model SVM dan CNN berdasarkan lima matrik evaluasi utama, ditampilkan dalam bentuk *bar chart*. Pada grafik ini terlihat bahwa model SVM menunjukkan performa lebih baik pada seluruh metrik, kecuali recall yang nilainya sama antara kedua model.



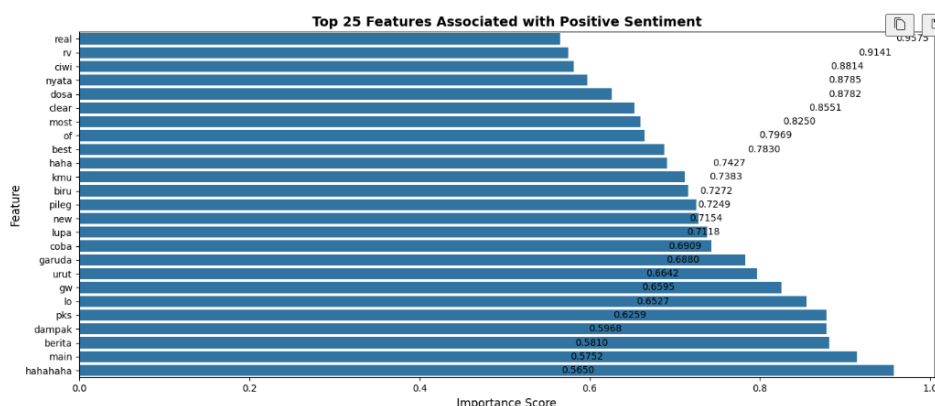
Gambar 20. Spider Chart Hasil Perbandingan

Gambar 20 menunjukkan radar *chart* atau *spider chart* yang digunakan untuk membandingkan performa dua model, yaitu SVM (garis biru) dan CNN (garis merah), berdasarkan beberapa matrik evaluasi model klasifikasi. Matrik yang ditampilkan berupa *precison*, *recall*, *acuracy*, *F1-score* dan *secificity*. Setiap sumbu mempresentasikan salah satu matrik evaluasi. Jarak dari pusat ke tepi menggambarkan nilai masing-masing matrik dari 0 hingga 1. Garis biru (SVM) menutupi area yang lebih luas, terutama pada ada matrik *accuracy*, *secificity* dan *recall*, menunjukkan performa yang lebih tinggi pada matrik tersebut. Garis merah (CNN) menutupi area yang lebih kecil, menandakan performa CNN yang lebih rendah dibanding SVM dalam matrik-matrik tersebut.



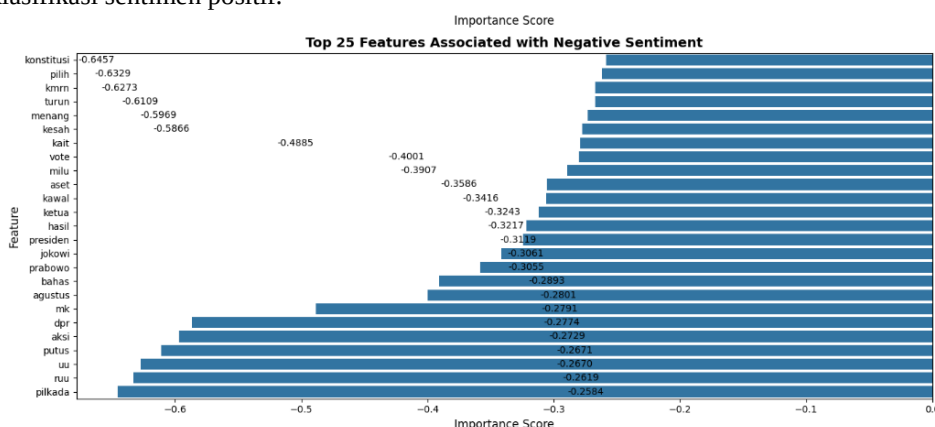
Gambar 21. Grafik Perbandingan Kinerja SVM dan CNN

Gambar 21 menunjukkan grafik perbandingan kinerja antara dua model SVM dan CNN berdasarkan nilai matrik *accuracy*, *precision*, *recall*, *secificity* dan *F1-score* dimana diagram berwarna biru mewakili model SVM dan diagram berwarna merah melambangkan model CNN. Dengan demikian dapat dilihat berdasarkan diagram diatas SVM mengungguli CNN secara signifikan dalam hampir semua matrik.



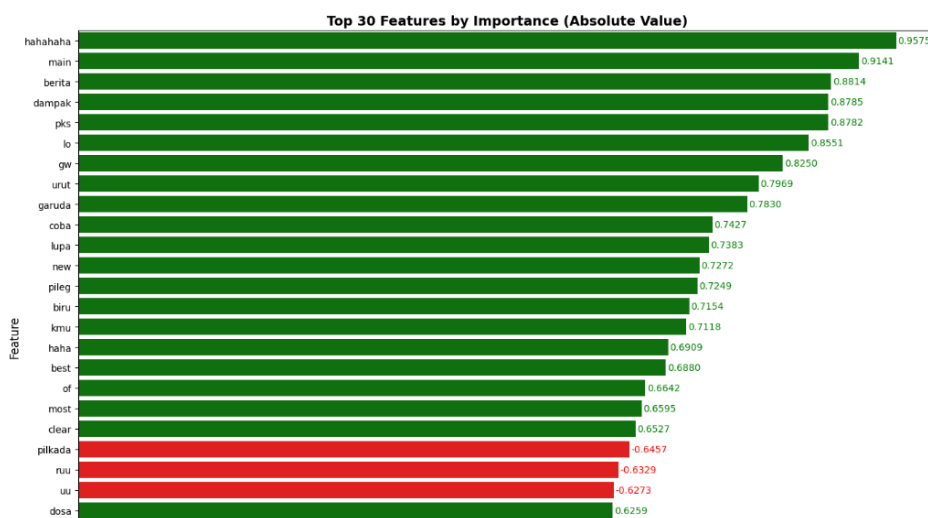
Gambar 22. Token Terpenting dalam Sentimen Positif

Gambar diatas menunjukkan 22 kata atau token yang paling berkontribusi terhadap prediksi sentimen positif. Sumbu Y berisi kata-kata (fitur) dari *tweet* yang dianggap sebagai indikator sentimen positif. Sumbu X adalah *importance score* atau skor penting fitur dalam menentukan sentimen positif. Semakin tinggi skornya, semakin besar pengaruh fitur tersebut dalam klasifikasi sentimen positif.



Gambar 23. Token Terenting dalam Sentimen Negatif

Gambar diatas menunjukkan 24 kata atau token yang paling berkontribusi terhadap prediksi sentimen negatif. Sumbu Y berisi kata-kata (fitur) dari *tweet* yang dianggap sebagai indikator sentimen negatif. Sumbu X adalah *importance score* atau skor penting fitur dalam menentukan sentimen negatif. Semakin tinggi skornya, semakin besar pengaruh fitur tersebut dalam klasifikasi sentimen negatif.

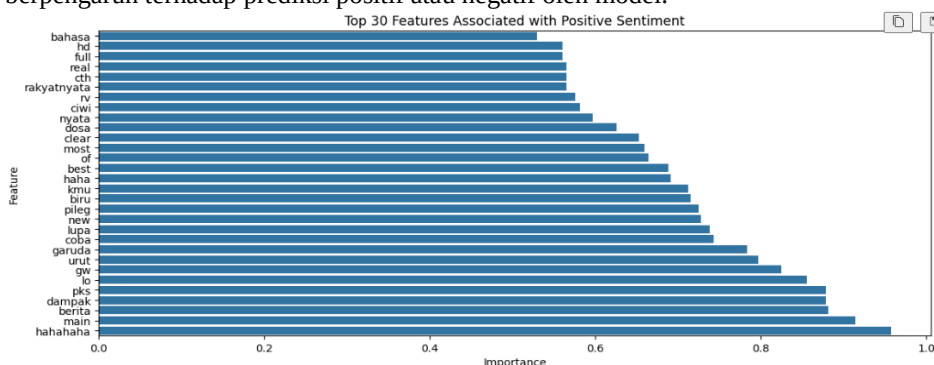


Gambar 24. Grafik Nilai Absolut

Grafik top 30 fitur berdasarkan nilai absolut yang ditunjukkan pada gambar 25 menampilkan kata-kata yang paling berpengaruh terhadap klasifikasi sentimen baik positif maupun negatif. Sumbu Y pada grafik berupa kata-kata fitur dari data sedangkan sumbu X merupakan *importance score* atau nilai kontribusi fitur terhadap keputusan model. Warna pada grafik juga memiliki makna dimana warna hijau melambangkan fitur yang berkorelasi dengan sentimen positif dan warna merah melambangkan fitur yang berkorelasi dengan sentimen negatif.

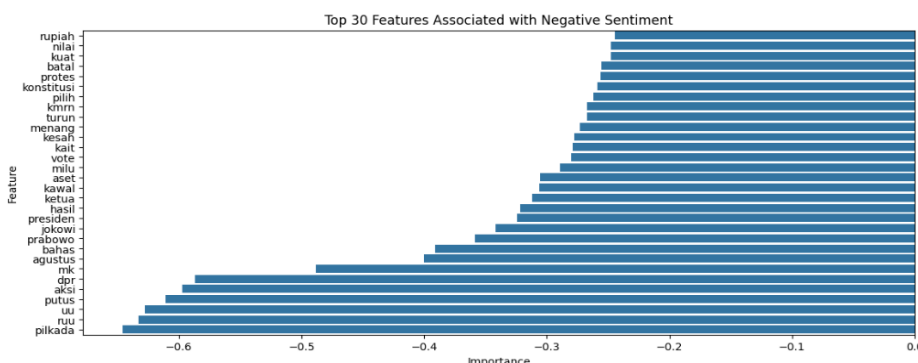
H. Analisis Feature Importance

Pada tahap ini dilakukan analisis pentingnya fitur dari model SVM dengan tujuan untuk mengetahui kata-kata atau fitur apa yang paling berpengaruh terhadap prediksi positif atau negatif oleh model.



Gambar 25. Hasil analisis fitur sentimen positif

Gambar 25 menunjukan 30 kata teratas yang paling berpengaruh terhadap sentimen positif berdasarkan analisis *feature importance* dari model.



Gambar 26. Hasil Analisis Fitur Sentimen Negatif

Gambar 26 menunjukkan 30 kata teratas yang paling berpengaruh terhadap sentimen negatif berdasarkan analisis *feature importance* dari model.

I. Laporan Hasil Analisis Sentimen RUU Pilkada dengan SVM dan CNN

Setelah melalui banyak proses maka dilakukan pembuatan matriks performa dari model SVM dan CNN, perhitungan skor rata-rata berbobot untuk menentukan model terbaik, analisis distribusi sentimen dari tweet, identifikasi kata-kata kunci yang berpengaruh, serta penyusunan laporan dan ringkasan eksekutif dalam format markdown untuk mendukung analisis dan pengambilan keputusan berbasis data.

```
Generating comprehensive analysis report...
SVM metrics tersedia dalam memory
CNN metrics tersedia dalam memory

SVM Metrics yang digunakan:
accuracy: 0.9238
precision: 0.2143
recall: 0.3750
specificity: 0.9455
f1_score: 0.2727

CNN Metrics yang digunakan:
accuracy: 0.6286
precision: 0.0395
recall: 0.3750
specificity: 0.6386
f1_score: 0.0714
Comprehensive analysis report generated successfully!

Analysis complete! All outputs are organized in the 'output' directory structure.
```

Gambar 27. Hasil Laporan Akhir

IV. KESIMPULAN

Penelitian ini membandingkan performa metode *Support Vector Machine (SVM)* dan *Convolutional Neural Network (CNN)* dalam analisis Sentimen terhadap RUU Pilkada di Aplikasi X. Berdasarkan hasil evaluasi, SVM menunjukkan kinerja yang lebih unggul dibandingkan CNN pada hampir semua matriks, dengan akurasi SVM mencapai 92,38% dan F1-score 27,27%. Sebaliknya, CNN mencatatkan akurasi 62,86 dan F1-score 7,14%. Meskipun CNN unggul dalam pemrosesan data sekuensial, SVM terbukti lebih stabil dan akurat dalam konteks dataset ini.

V. DAFTAR PUSTAKA

- [1] Rianda, M., dan Hamzah, M. A. Implikasi Putusan Mahkamah Konstitusi Nomor 1/PHPU. PRES-XXII/2024 terhadap Pelaksanaan Pemilihan Kepala Daerah.
- [2] Tokan, C. F., Andriani, P., Ang, S. Dan Prianto, Y. 2024. Penanganan Demonstran Pada Aksi

- Mengawal Putusan Mk Di Jakarta. *Multilingual: Journal Of Universal Studies*. 4(3): 368-380, available at: <https://ejournal.penerbitjurnal.com/index.php/multilingual/article/view/928>
- [3] Nursinggah, L., Mufizar, T. Dan Perjuangan, U. 2024. Analisis Sentimen Pengguna Aplikasi X Terhadap Program Makan Siang Gratis Dengan Metode Naïve Bayes Classifier. *J. Inform. Dan Tek. Elektro Ter.* 12(3). Available at: <https://journal.eng.unila.ac.id/index.php/jitet/article/view/4336>
- [4] Naufal, M. F. Dan Kusuma, S. F. 2022. Analisis Sentimen Pada Media Sosial Twitter Terhadap Kebijakan Pemberlakuan Pembatasan Kegiatan Masyarakat Berbasis Deep Learning. *Jurnal Edukasi Dan Penelitian*. 8(1): 44-49. Available at: <https://repository.ubaya.ac.id/41763/>
- [5] Sitio, G. Y., Rumapea, S. A. Dan Lumbanraja, P. 2023. Analisis Sentimen Pemindahan Ibu Kota Negara Di Media Sosial Twitter Menggunakan Metode Convolutional Neural Network (CNN). *METHOTIKA: Jurnal Ilmiah Teknik Informatika*. 3(2): 97-104. Available at: <https://ejurnal.methodist.ac.id/index.php/methotika/article/view/2680>
- [6] Susilawati, A. T., Lestari, N. A. Dan Nina, P. A. 2024. Analisis Sentimen Publik Pada Twitter Terhadap Boikot Produk Israel Menggunakan Metode Naïve Bayes. *Nian Tana Sikka: Jurnal Ilmiah Mahasiswa*. 2(1): 26-35. Tersedia pada: <https://ejournal-nipamof.id/index.php/NianTanaSikka/article/view/240>
- [7] Abidin, R. Dan Herawati, A. 2024. Analisis Sentimen Publik Terhadap Kebijakan Program Tabungan Perumahan Rakyat (Tapera). *Journal Of Information System And Computer*. 4(1): 13-19. available at: <https://journal.unisnu.ac.id/JISTER/article/view/1002/449>
- [8] Setyaningtyas, E. Dan Nugroho, K. 2024. Analisis Sentimen Media Sosial Pada Pengguna Twitter Terhadap Pemilu 2024 Menggunakan Metode LSTM. *Jurasik (Jurnal Riset Sistem Informasi Dan Teknik Informatika)*. 9(2): 673-683. available at: <https://tunasbangsa.ac.id/ejurnal/index.php/jurasik/article/view/799>
- [9] Ramdhani, N. 2021. Analisis Sentimen Pengguna Twitter Terhadap Belajar Daring Selama Pandemi Covid-19 Dengan Deep Learning. *Jurnal Siliwangi Seri Sains Dan Teknologi*. 7(2). available at: <https://jurnal.unsil.ac.id/index.php/jssainstek/article/view/4281>
- [10] Harahap, E. D. Dan Kurniawan, R. 2024. Analisis Sentimen Komentar Terhadap Kebijakan Pemerintah Mengenai Tabungan Perumahan Rakyat (TAPERA) Pada Aplikasi X Menggunakan Metode Naïve Bayes. *Jurnal Teknik Informatika UNIKA Santo Thomas*. 166-175. available at: <https://ejournal.ust.ac.id/index.php/JTIUST/article/view/3911>
- [11] Maulana, A. Dan Yuliana, A. 2024. Analisis Sentimen Opini Publik Terkait Judi Online Pada Pengguna Aplikasi X Menggunakan Algoritma Naïve Bayes Dan Support Vector Mechine. *Jurnal Informatika Dan Teknik Elektro Terapan*. 12(3S1): available at: <https://journal.eng.unila.ac.id/index.php/jitet/article/view/5187>
- [12] Wijayanti, K. Dan Wijayani, Q. N. 2024. Peranan Aplikasi Twitter Atau X Dalam Interaksi Komunikasi Guna Membantu Penyeimbangan Kesehatan Mental Pada Remaja Saat Ini. *Journal Sains Student Research*. 2(1): 07-15. available at: <https://ejurnal.kampusakademik.co.id/index.php/jssr/article/view/469>
- [13] Homaidi, A. Dan Fatah, Z. 2024. Implementasi Metode K-Nearest Neighbors (KNN) Untuk Klasifikasi Penyakit Jantung. *G-Tech: Jurnal Teknologi Terapan*. 8(3): 1720-1728. Available at: <https://ejournal.uniramalang.ac.id/index.php/g-tech/article/view/4495>
- [14] Purnama, J. J. Dan Rahayu, S. 2022. Klasifikasi Konsumsi Energi Industri Baja Menggunakan Teknik Data Mining. *Jurnal Teknoinfo*. 16(2): 395-407. available at: <https://ejurnal.teknokrat.ac.id/index.php/teknoinfo/article/view/1984>
- [15] Paul, H., Wiguna, A. S. Dan Santoso, H. 2023. Penerapan Algoritma Support Vector Machine Dan Naive Bayes Untuk Klasifikasi Jenis Mobil Terlaris Berdasarkan Produksi Di Indonesia. *JATI (Jurnal Mahasiswa Teknik Informatika)*. 7(1): 39-44. available at: <https://journal.eng.unila.ac.id/index.php/jitet/article/view/4336>
- [16] Susana, H. 2022. Penerapan Model Klasifikasi Metode Naive Bayes Terhadap Penggunaan Akses Internet. *Jurnal Riset Sistem Informasi Dan Teknologi Informasi (JURSISTEKNI)*. 4(1): 1-8. available at: <https://jursistekni.nusaputra.ac.id/article/view/96>
- [17] Larasati, R. C., Dewi, C. Dan Christanto, H. J. 2024. Analisis Sentimen Produk Kecantikan



- Jenis Moisturizer Di Twitter Menggunakan Algoritma Super Vector Machine. *Jurnal Tekinkom (Teknik Informasi Dan Komputer)*. 7(1): 124-134. available at: <https://jurnal.murnisadar.ac.id/index.php/Tekinkom/article/view/1243>
- [18] Maulana, A. R. Dan Rochmawati, N. 2020. Opinion Mining Terhadap Pemberitaan Corona Di Instagram Menggunakan Convolutional Neural Network. *Journal Of Informatics And Computer Science (JINACS)*. 2(01): 53-59. available at: <https://ejournal.unesa.ac.id>
- [19] Permana, M. A., Darwiyanto, E. Dan Bijaksana, M. A. 2021. Pembobotan Dan Pemeringkatan Ayat Al-Quran Berdasarkan Compound, Term Frequency Dan Prinsip Pareto Untuk Membantu Hafalan. *Eproceedings Of Engineering*. 8(2). available at: <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/14723>
- [20] Rohmansa, R. Q., Pratiwi, N. Dan Palepa, M. J. 2024. Analisis Sentimen Ulasan Pengguna Aplikasi Discord Menggunakan Metode K-Nearest Neighbor. *JUPI (Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika)*. 9(1): 368-378. available at: <https://jurnal.stkipgritulungagung.ac.id/index.php/jupi/article/view/4943>
- [21] Liawati, A., Narasati, R., Solihudin, D., Rohmat, C. L. Dan Permana, S. E. 2023. Analisis Sentimen Komentar Politik Di Media Sosial X Dengan Pendekatan Deep Learning. *JATI (Jurnal Mahasiswa Teknik Informatika)*. 7(6): 3557-3563. available at: <https://ejournal.itn.ac.id/index.php/jati/article/view/8248>
- [22] Arfan, I. S., Fauziah, S. Dan Nawangsih, I. 2024. Analisis Sentimen Terhadap Cyber Bullying Di X Menggunakan Algoritma Naïve Bayes: Sentiment Analyst Of Cyber Bullying In X Using Naïve Bayes Algorithm. *MALCOM: Indonesian Journal Of Machine Learning And Computer Science*. 4(4): 1411-1419. available at: <https://journal.irpi.or.id/index.php/malcom>
- [23] Ni'mah, A. T. Dan Arifin, A. Z. 2020. Perbandingan Metode Term Weighting Terhadap Hasil Klasifikasi Teks Pada Dataset Terjemahan Kitab Hadis. 13(2): 172-180. available at: <https://journal.trunojoyo.ac.id/rekayasa/article/view/6412>
- [24] Huda, K., Pohan, S. D., dan Herlina, Y. 2024. Penerapan Pembobotan Term Frequency InverseDocument Frequency dan Algoritma K-Nearest Neighbor untuk Analisis Ulasan Hotel di Situs Tripadvisor. *Jurnal Informatika dan Teknik Elektro Terapan*. 12(3). available at: <https://journal.eng.unila.ac.id/index.php/jitet/article/view/4800>
- [25] Lestari, S. Dan Febryanti, M. 2024. Analisis Sentimen Mengenai Produk Inovasi Invisible TV Menggunakan Algoritma Naïve Bayes. *INTECOMS: Journal Of Information Technology And Computer Science*. (5): 1675-1684. available at: <https://journal.ipm2kpe.or.id/index.php/INTECOM/article/view/11585/7824>
- [26] Yohannes, R. Dan Al Rivan, M. E. 2022. Klasifikasi Jenis Kanker Kulit Menggunakan CNN SVM. *Jurnal Algoritme*. 2(2): 133-144. available at: <https://jurnal.mdp.ac.id/index.php/algoritme/article/view/2363>
- [27] Aruan, J. D. C., Rahayudi, B. Dan Ridok, A. 2022. Analisis Sentimen Opini Masyarakat Terhadap Pelayanan Rumah Sakit Umum Daerah Menggunakan Metode Support Vector Machine Dan Term Frequency-Inverse Document Frequency. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*. 6(5): 2072-2078. available at: <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/10976/4860>
- [28] Adek, R. T., Fikry, M. Dan Khalil, U. 2021. News Opinion Classification Application With Support Vector Machine Algorithm Using Framework Codeigniter. *Journal Of Informatics And Telecommunication Engineering*. 5(1): 160-166. available at: <https://ojs.uma.ac.id/index.php/jite/article/view/5189/3420>
- [29] Sari, D. Dan Kurniawan, R. 2024. Analisis Sentimen Terhadap Kinerja Program Walikota Medan Pada Media Sosial X Menggunakan Support Vector Machine: Sentiment Analysis Of The Performance Of Medan Mayor Program On Social Media X Using Support Vector Machine. *MALCOM: Indonesian Journal Of Machine Learning And Computer Science*. 4(4): 1539-1548. available at: <https://journal.irpi.or.id/index.php/malcom/article/view/1685>
- [30] Adiyatma, F. A., Alam, S. Dan Komara, M. A. 2024. Analisis Sentimen Masyarakat Di Platform X Terhadap

Penggunaan Bansos Untuk Memenangkan Salah Satu Capres Tertentu Di Pilpres 2024 Menggunakan Metode Naïve Bayes Classifier. *Jati (Jurnal Mahasiswa Teknik Informatika)*. 8(5): 9941-9947. available at: <https://ejournal.itn.ac.id/index.php/jati/article/view/10836>

- [31] Wati, R. Dan Ernawati, S. 2021. Analisis Sentimen Persepsi Publik Mengenai PPKM Pada Twitter Berbasis SVM Menggunakan Python. *Jurnal Teknik Informatika UNIKA Santo Thomas*. 240-247. available at: <https://ejournal.ust.ac.id/index.php/JTIUST/article/view/1465>
- [32] Vindua, R. Dan Zailani, A. U. 2023. Analisis Sentimen Pemilu Indonesia Tahun 2024 Dari Media Sosial Twitter Menggunakan Python. *JURIKOM (Jurnal Riset Komputer)*. 10(2): 479-487. available at: <https://ejurnal.stmik-budidarma.ac.id>



DOI: 10.52362/jisamar.v9i2.1887

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).