

SOCIAL NETWORK ANALYSIS PERINGATAN DARURAT RUU PEMILIHAN KEPALA DAERAH

Stacia Christiano, Jusia Amanda Ginting
Teknik Informatika, Universitas Bunda Mulia
Jl. Lodan Raya Ancol, Jakarta Utara, Indonesia
S32210044@student.ubm.ac.id

ABSTRAK

Tagar #KawalPutusanMK menjadi trending topic pada 21 Agustus 2024, menimbulkan beragam reaksi masyarakat terhadap putusan Mahkamah Konstitusi terkait perubahan rancangan peraturan tersebut. Penelitian ini bertujuan untuk memetakan sebaran informasi menggunakan *Social Network Analysis* (SNA) dan analisis sentimen publik menggunakan algoritma *Support Vector Machine* (SVM). Metodologi penelitian ini menggunakan pendekatan kuantitatif berbasis data. Data dikumpulkan di media sosial X dengan tagar #KawalPutusanMK. Analisis dilakukan dengan SNA untuk mengidentifikasi aktor kunci dalam jaringan sosial berdasarkan *degree centrality*, *betweenness centrality* dan *closeness centrality*. Untuk analisis sentimen, teknik TF-IDF digunakan untuk mengekstraksi fitur dan algoritma SVM diterapkan untuk mengklasifikasikan sentimen ke dalam tiga kategori. Hasil penelitian menunjukkan bahwa aktor utama yang berperan penting dalam penyebaran informasi adalah dessertmeys, dimm_sky, zyze9, hasbil_lbs, lalalabloem, dsperdana, ddazling_ dan mamtaltakiri. Dalam analisis sentimen, kernel linier memberikan akurasi tertinggi sebesar 67%, diikuti oleh sigmoid sebesar 64%, RBF sebesar 63%, dan polinomial sebesar 59%. Opini masyarakat sebagian besar bersifat negatif (51,18%), sedangkan tanggapan netral dan positif masing-masing mencapai 26,38% dan 22,44%. Penelitian ini menyimpulkan bahwa media sosial X adalah platform penting untuk penyebaran informasi dan analisis opini publik. Saran dari penelitian ini adalah meningkatkan metode pelabelan data dan memilih algoritma yang lebih tepat untuk meningkatkan akurasi analisis di masa depan.

Kata kunci : Analisis Sentimen, Media Sosial, *Social Network Analysis*, *Support Vector Machine*

1. PENDAHULUAN

Media sosial merupakan media online yang digunakan untuk dapat terhubung tanpa adanya batasan ruang dan waktu. Bentuk komunikasi tradisional semakin tergerus berkat adanya media sosial yang menyediakan efektifitas komunikasi dan penyebaran informasi secara cepat kepada publik [1][2].

Mengutip dari laman we are social tahun 2023, pengguna X di Indonesia mencapai 60,4 %. Keunggulan media sosial ini dibandingkan dengan media sosial yang lain adalah kemudahan pengguna untuk saling berinteraksi dan berkomunikasi dengan pengguna lain (Putra, 2014) [2].

Tagar #KawalPutusanMK menjadi trending topik di media sosial X pada tanggal 21 Agustus 2024 terkait dengan keputusan Mahkamah Konstitusi (MK) terhadap ketentuan perubahan regulasi rancangan undang-undang. Tentu saja perubahan ini menjadi pusat perhatian masyarakat Indonesia. Banyak masyarakat memberikan komentar positif, negatif, dan netral.

Penelitian ini memiliki tujuan untuk melihat sebaran informasi dari tagar tersebut dengan menggunakan SNA. SNA adalah sebuah metode yang memvisualisasikan graf melalui hubungan antar manusia[3].

Selain itu, penelitian ini juga akan mengklasifikasikan postingan positif, netral, dan negatif guna mengetahui keseluruhan sentimen dalam

penyebaran #KawalPutusanMK di media sosial X dengan menggunakan Algoritma SVM.

SVM sendiri memiliki fungsi untuk mengklasifikasikan sentimen baik dalam data linear maupun non linear. Penelitian ini menggunakan SVM sebagai algoritma untuk analisis sentimen karena pada penelitian [12] hasil akurasi SVM lebih tinggi dibandingkan dengan Naïve Bayes.

Penelitian yang telah dilakukan akan menghasilkan 10 aktor kunci yang berpengaruh besar dalam penyebaran postingan tersebut, mendapatkan hasil dari klasifikasi sentimen yaitu positif, negatif dan netral, serta mendapatkan perbandingan akurasi menggunakan empat kernel SVM.

2. TINJAUAN PUSTAKA

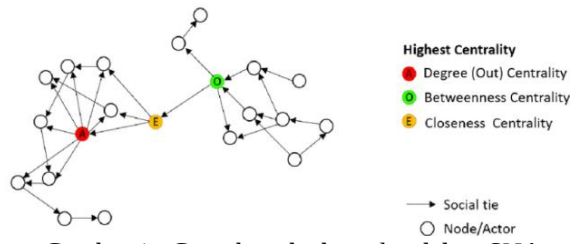
2.1. Media Sosial X sebagai Sumber Informasi

Media sosial merupakan media yang paling efektif dalam penyebaran informasi kepada publik. Dikarenakan informasi yang tersebar melalui media sosial tidak perlu didistribusikan lagi ke publik secara fisik, cukup hanya dengan memiliki akses internet[1].

2.2. *Social Network Analysis* (SNA)

Social Network Analysis (SNA) adalah sebuah metode untuk mempelajari hubungan antar manusia dengan menggunakan teori graf [3]. SNA dapat digunakan untuk mempelajari pola jaringan organisasi, ide, dan orang yang terhubung dengan berbagai cara [5]. Dalam SNA, individu atau objek digambarkan sebagai node atau titik, sedangkan relasi

yang terjadi antar individu atau objek disebut edge atau link. Secara umum, sebuah jaringan sosial adalah sebuah peta yang terdiri dari banyak objek yang memiliki relasi antar objek-objek itu sendiri [3].



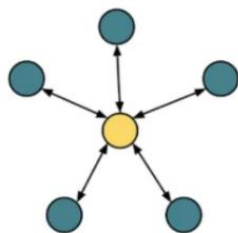
Gambar 1. Contoh node dan edge dalam SNA

Gambar 1 adalah contoh graf yang meliputi node dan edge yang dipresentasikan menggunakan tiga metrik *centrality measures*.

Centrality measures terdiri dari :

2.3. Degree Centrality

Degree centrality digunakan untuk menentukan jumlah dari edge yang terhubung ke node yang memiliki nilai degree centrality dan berbeda dari node lainnya [3]. *Degree Centrality* menghitung jumlah koneksi atau interaksi yang dimiliki oleh sebuah node [4][10].



Gambar 2. Graf degree centrality

Rumus *Degree Centrality* :

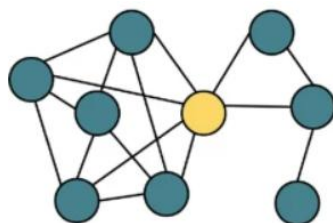
$$Cd(v) = \frac{dv}{n-1} \quad (1)$$

Di mana :

$Cd(v)$ = degree centrality node v
 dv = jumlah edge yang dimiliki node v
 n = jumlah node dalam jaringan

2.4. Betweenness Centrality

Betweenness Centrality menghitung seberapa sering sebuah node dilewati dengan node lain untuk menuju ke node tertentu dalam jaringan. Nilai ini berfungsi untuk menentukan peran aktor yang menjadi jembatan yang menghubungkan interaksi dalam jaringan [3][4].



Gambar 3. Graf betweenness centrality

Rumus *Betweenness Centrality*:

$$Cb(v) = \sum_{j=1}^N \sum_{k=1}^{j-1} \frac{gjk(i)}{gjk} \quad (2)$$

Di mana :

$Cb(v)$: Betweenness centrality dari node v

N : Jumlah total node dalam graf

$gjk(i)$: Jumlah jalur terpendek dari node j ke node k yang melewati node i.

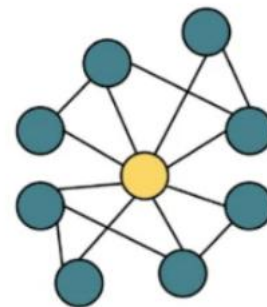
gjk : Jumlah jalur terpendek dari node j ke node k yang melewati node i atau tidak.

$\sum_{j=1}^N$: Menyatakan jumlah untuk setiap pasangan node j dan k, di mana j dan k adalah node yang berbeda dalam graf.

$\sum_{k=1}^{j-1}$: Menunjukkan bahwa kita mempertimbangkan setiap pasangan node j dan k yang berbeda, di mana k lebih kecil dari j untuk menghindari perhitungan duplikat.

2.5. Closeness Centrality

Closeness centrality adalah rata-rata dari keseluruhan jalur terpendek dari satu node ke semua node lain di jaringan. *Closeness centrality* menghitung jarak rata-rata antara sebuah node dan semua node lain di jaringan [3][4].



Gambar 4. Graf closeness centrality

Rumus *Closeness Centrality* :

$$Cc(i) = \frac{n-1}{\sum dij} \quad (3)$$

Di mana :

$Cc(v)$ = Closeness centrality node v

dij = jarak antara node i dan j

n = jumlah node dalam jaringan

2.6. Analisis Sentimen

Analisis sentimen atau opinion mining merupakan proses memahami, mengekstrak dan mengolah data tekstual secara otomatis untuk mendapatkan informasi sentimen yang terkandung dalam suatu kalimat opini [5]. Analisis sentimen dipergunakan untuk melihat pendapat atau kesamaan opini terhadap sebuah persoalan atau objek oleh seorang menuju ke opini positif atau negatif [6][7]. Analisis sentimen mempunyai beberapa tahapan agar dapat menghasilkan *output* yang sesuai, yaitu :

2.7. Tahap Pre Processing

Pre-processing merupakan tahapan mengolah data berupa teks dengan cara menghilangkan data-data yang tidak diperlukan untuk proses selanjutnya agar bisa dilakukan analisis sentimen (Widagdo et al., 2020) [6]. Hasil dari proses *pre-processing* memiliki dampak besar terhadap data yang diolah dan bisa memengaruhi akurasi klasifikasi dokumen[8]. Berikut adalah tahap *pre-processing* data :

a. Cleansing

Cleansing merupakan hasil *scraping* data dari Twitter, tahap *cleansing* adalah tahap eliminasi aksara. Aksara yang dihapus adalah tanda baca seperti titik (.), koma (,), tanda tanya (?), dan tanda seru (!), serta simbol-simbol seperti tanda '@' untuk *username*, tagar (#), *emoticon*, dan alamat *website*[7][13].

b. Case Folding

Case folding dapat mengubah bentuk kata menjadi bentuk dasarnya agar sebuah karakter dapat seragam (*lower case*) atau dapat disebut juga penyeragaman bentuk huruf [6].

c. Tokenizing

Tokenizing dapat memisahkan data teks menjadi beberapa token. Pada prinsipnya proses *Tokenizing* adalah memisahkan setiap kata yang menyusun suatu dokumen. Pada umumnya setiap kata teridentifikasi atau terpisahkan dengan kata yang lain oleh karakter spasi, sehingga proses *tokenizing* mengandalkan karakter spasi pada dokumen untuk melakukan pemisahan kata [5][6].

d. Filtering

Filtering atau *stop removal* merupakan proses membuang kata yang tidak penting dari proses *tokenizing* sebelumnya[6]. Kata-kata yang dimaksud seperti “yang”, “di”, “dan”, “pada”, dan sebagainya[9].

e. Stemming

Stemming yaitu melakukan proses mencari kata dasar dari setiap kata hasil proses *filtering* sebelumnya[6]. Tahap ini berfungsi mengubah kata berimbuhan pada tiap kata yang telah terseleksi menjadi kata dasar[19].

f. Normalization

Normalization dilakukan untuk mengubah kata dari yang tidak baku menjadi baku, tahap ini dilakukan menggunakan database kamus kata bahasa baku dan tidak baku yang dibuat sendiri berdasarkan dari data komentar yang digunakan[6].

2.8. Labeling menggunakan Lexicon Inset

Lexicon Inset adalah pengembangan kamus *lexicon* dengan bahasa indonesia dengan tujuan untuk menganalisa sentimen. *Lexicon Inset* dibuat pada tahun 2018 dengan menggunakan data dari X dengan isi kamus 3609 kata positif dan 6609 negatif. Pembobotan kata *Lexicon Inset* diantara -5 sampai 5 yang diberikan secara manual. Penentuan jenis sentimen dilakukan dengan total bobot pada sebuah

kata, jika total bobot lebih dari 0 maka masuk kedalam label positif, jika total bobot kurang dari 0 maka masuk kedalam sentimen negatif, dan apabila bobot sama dengan 0 maka termasuk dalam sentimen netral [10].

2.9. Pembobotan TF IDF

Term Frequency (TF) merupakan jumlah kemunculan atau frekuensi kata pada suatu dokumen. Sedangkan *Inverse Document Frequency* (IDF) bertujuan untuk mengetahui apakah *term* yang dicari cocok dengan kata kunci yang diinginkan, *term* yang sering muncul akan memberikan pengaruh yang kecil dalam menentukan keterkaitan kata kunci dokumen[5].

Berikut persamaan TF IDF:

$$TF - IDF(d, t) = TF(d, t) * IDF(t) \quad (4)$$

Perhitungan TF-IDF bersumber dari perhitungan TF dan IDF menggunakan rumus berikut

$$TF(d, t) = \frac{\text{Jumlah kata } t \text{ pada dokumen } d}{\text{total kata pada dokumen } d} \quad (5)$$

$$IDF(t) = \frac{1}{\log \frac{\text{total dokumen}}{\text{Jumlah dokumen yang mengandung kata } t}} \quad (6)$$

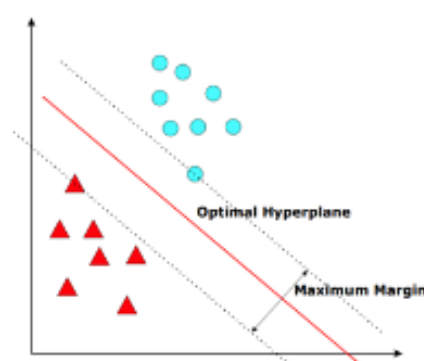
Dimana :

t = kata

d = dokumen

2.10. Klasifikasi SVM

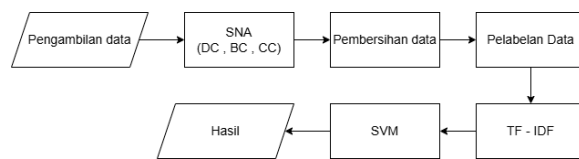
Support Vector Machine (SVM) adalah algoritma *supervised-learning* yang sering digunakan dalam klasifikasi dan regresi. SVM bekerja dengan cara mencari *hyperplane* yang optimal. *Hyperplane* adalah pemisah antara dua kelas yang fungsinya membagi sampel menjadi dua. *Hyperplane* yang memisahkan sampel memiliki jarak *margin* maksimum (*maximum-margin hyperplane*) yang dihitung dari *hyperplane* ke vektor terdekat [11].



Gambar 5. Maksimum margin memisahkan dua kelas

Terdapat dua jenis *hyperplane* dalam SVM yaitu linear dan non linear. Jika data dapat dipisahkan dengan sempurna oleh sebuah *hyperplane* linear maka SVM disebut dengan SVM linear.

3. METODE PENELITIAN



Gambar 6. Tahapan penelitian

Berikut merupakan langkah-langkah yang dilakukan berdasarkan tahapan penelitian pada gambar 6.

3.1. Analisis Kebutuhan

Tahap analisis kebutuhan adalah tahap awal di mana peneliti menentukan tujuan penelitian dengan cara mengidentifikasi topik permasalahan yang akan di angkat , sumber data , perangkat yang digunakan dan metode atau algoritma yang digunakan dalam penelitian ini.

Data yang dianalisis bertajuk “Peringatan Darurat” dikumpulkan dari media sosial X (Twitter) dengan menggunakan tagar #KawalPutusanMK. Data yang dikumpulkan dari rentang waktu satu bulan , sehingga data yang dikumpulkan akan lebih spesifik , dikarenakan kemunculan tagar #KawalPutusanMK ini sejak 21 Agustus 2024.

NodeXL adalah alat analisis jejaring sosial yang dapat dihubungkan ke Microsoft Excel. Berfungsi membantu penelitian ini untuk mengimpor, menganalisis, dan memvisualisasikan data.

Dataset awal yang telah dikumpulkan melalui proses *crawling* menjadi dasar penelitian ini untuk dilanjutkan ke tahapan selanjutnya yaitu melakukan analisis sentimen. Sebelum melakukan analisis , dataset awal akan ditransformasikan melalui tahapan *pre processing* dan *labeling* untuk membersihkan data sehingga data siap di analisis.

3.2. Desain

Desain dilakukan untuk mempersiapkan data hingga siap di analisis. Desain memiliki beberapa tahap yang harus diperhatikan. Tahap pertama penelitian ini melakukan *pre processing* untuk membersihkan dan mempersiapkan data agar siap dianalisis menggunakan *Social Network Analysis* (SNA) dan Analisis Sentimen. Berikut adalah tahapan desain.

3.3. Pre Processing

Pre processing dilakukan di Google Colab menggunakan *python* sebagai bahasa pemrograman.

a. Cleansing Data

Cleansing dilakukan untuk membersihkan URL , tagar atau hashtag (#) , username (@username) , email , emoticon (:@ , :* , :D), tanda baca seperti koma (,) , titik (.) dan juga tanda baca lainnya. *Case folding* juga termasuk dalam tahap *cleansing* data. *Case folding* dilakukan untuk mengubah setiap

data yang ada dalam dokumen menjadi huruf kecil menggunakan (*lower case*).

b. Tokenizing

Tahap *tokenizing* dalam penelitian ini dilakukan untuk memecah teks menjadi token atau *term*. Token disini merupakan kata yang dipecah berdasarkan spasi.

c. Filtering

Filtering atau *stop removal* adalah tahapan untuk menghapus kata yang sering muncul namun tidak penting. Seperti kata “ada” , “yang” , “dan” , “bisa” dan banyak kata lainnya.

d. Stemming

Tahap *stemming* dilakukan untuk menjadikan sebuah kata menjadi kata dasar dengan menghapus imbuhan, awalan, dan akhiran.

e. Normalization

Tahap *normalization* dalam penelitian ini dilakukan untuk mengubah kata yang tidak baku menjadi baku , serta menjadikan kata yang sudah di *stemmed* menjadi satu kalimat

3.4. Labeling dengan Lexicon Inset

Labeling dengan *Lexicon Inset* dilakukan untuk memberikan kategori sentimen (positif , negatif, atau netral) pada data teks yang sudah dibersihkan. *Lexicon* adalah kamus kata yang sudah memiliki skor sentimen , kamus yang digunakan bersumber dari *GitHub*

(<https://github.com/fajri91/InSet>).

3.5. Implementasi

Tahap implementasi meliputi visualisasi graf serta perhitungan *centrality measures* dalam SNA dan analisis sentimen yang dilakukan menggunakan pembobotan TF-IDF dan klasifikasi menggunakan SVM.

a. Social Network Analysis (SNA)

SNA berfokus pada perhitungan *centrality measures* seperti *degree centrality*, *betweenness centrality* dan *closeness centrality*.

b. Pembobotan TF – IDF

Pembobotan TF-IDF dilakukan untuk mengubah data teks (*tweet*) menjadi representasi numerik (*vectorized form*) yang dapat diproses oleh algoritma SVM. TF-IDF dilakukan untuk menentukan pentingnya setiap kata dalam dokumen berdasarkan frekuensi kata tersebut dalam dokumen tertentu (TF) dan keberadaannya di seluruh dokumen (IDF).

c. Klasifikasi SVM

Support Vector Machine (SVM) dalam penelitian ini berfungsi untuk mengklasifikasikan tiga kategori sentimen, yaitu positif, negatif, dan netral. Langkah utama untuk menjalankan SVM adalah mengetahui kernel apa yang cocok dengan data yang digunakan.

3.6. Pengujian

Confusion Matrix merupakan alat untuk mengevaluasi kinerja model klasifikasi data *mining*.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (7)$$

$$Precision = \frac{TP}{TP + FP} \quad (8)$$

$$Recall = \frac{TP}{TP + FN} \quad (9)$$

$$F-1 \text{ Score} = 2 \times \frac{Recall \times Precision}{Recall + Precision} \quad (10)$$

Keterangan:

- Akurasi = perbandingan jumlah prediksi yang benar
- Presisi = proporsi kasus dengan hasil positif yang benar
- True Positive* (TP) = jumlah record positif yang diklasifikasikan oleh *classifier* sebagai positif.
- True Negative* (TN) = jumlah record negatif yang diklasifikasikan sebagai negatif oleh *classifier*.
- False Positive* (FP) = : jumlah record negatif yang dianggap positif oleh *classifier*.
- False Negative* (FN) = jumlah record positif yang diklasifikasikan sebagai negatif oleh *classifier*.

4. HASIL DAN PEMBAHASAN

4.1. Tahap Pengumpulan Data

Vertex 1	Vertex 2	Tweet
0 sliceofangelina	dessertmeys	@ckelatkyjy #KawalPutusanMK ln#TolakPilkadaAk...
1 weewooings	weewooings	Ada yang silent majority pas kawalputusanMK ke...
2 weewooings	adysauruss	Temen gue ada udah miskin, ngatain temen gue y...
3 cloveradishi	dessertmeys	@ckelatkyjy #KawalPutusanMK ln#TolakPilkadaAk...
4 hanyahanin	hanyahanin	Ketum Rampai Nusantara: Aksi Kawal Putusan MK ...
5 hanyahanin	hanyahanin	NasionalAksi Kawal Putusan MK Ada Indikasi Dir...
6 hanyahanin	hanyahanin	Penanganan Demo Kawal Putusan MK oleh Polri Su...
7 hanyahanin	hanyahanin	https://t.co/SCUeuO4P5F Rampai Nusantara Pasti...
8 hanyahanin	hanyahanin	Soal Pengamanan di Demo Kawal Putusan MK, Polr...
9 hanyahanin	hanyahanin	PB SEMMI MP, Bobby Kurniawan: Apresiasi Polri ...

Gambar 7. Kolom yang digunakan

Dari hasil *crawling*, penelitian ini menggunakan tiga kolom untuk dapat melakukan tahapan analisis. Kolom yang digunakan untuk penelitian ini yaitu *Vertex 1*, *Vertex 2* dan *Tweet* seperti pada gambar 7.

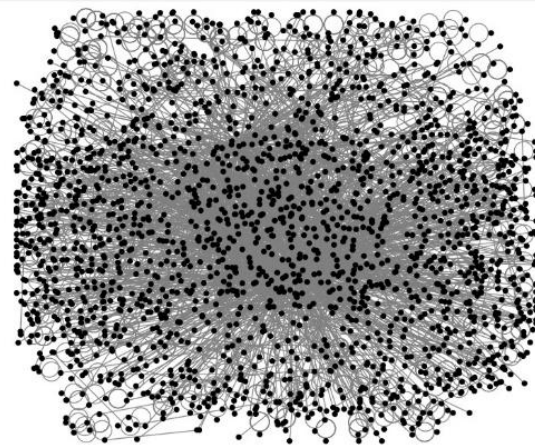
4.2. Social Network Analysis (SNA)

Graph yang akan digunakan dalam penelitian ini adalah *undirected graph* seperti pada gambar 8. Karena jika menggunakan *Directed Graph*, maka perhitungan *betweenness* sulit dihitung dikarenakan node terhubung menggunakan *degree centrality*.

4.3. Tahap Pre-processing

Tabel 1. Data hasil *pre processing*

Tweet	Normalized
Ada yang <i>silent majority</i> pas kawalputusanMK kemarin, terus ikut daftar cpns baru mencak-mencak peruri korupsi gegara udah beli e-materai tapi gak masuk-masuk ckckck	diam mayoritas ketika kawal putus mahkamah konstitusi kemarin daftar cpns mencak curi korupsi akibat sudah beli materai tidak masuk

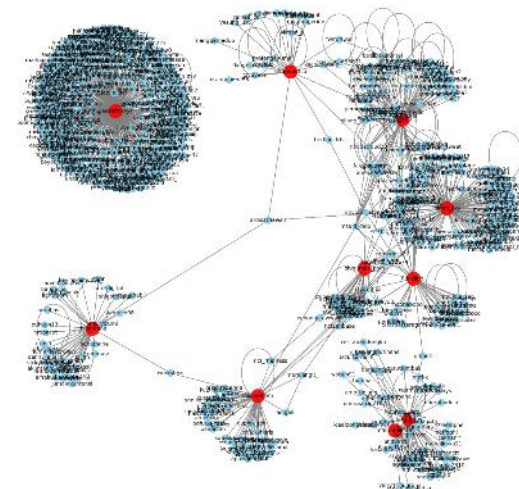


Gambar 8. Undirected graph

10 Teratas Aktor Kunci :

actor	degree	betweenness	closeness
1 dessertmeys	0.171442	0.182965	0.067810
106 dimm_sky	0.054881	0.101989	0.104265
228 zyze9	0.033997	0.069755	0.102612
294 hasbil_lbs	0.025255	0.041016	0.077968
355 lalalabloem	0.023798	0.191942	0.110439
1685 dsperdana	0.018941	0.024122	0.090372
1149 ddazling_	0.018456	0.074521	0.100493
209 mamatakatiri	0.012627	0.009562	0.069569
291 tempodotco	0.012142	0.017194	0.091844
176 divhumas_polri	0.011656	0.018070	0.097455

Gambar 9. Aktor kunci



Gambar 10. Graf 10 aktor kunci

Setelah mendapatkan data pengguna (*Vertex 1* dan *Vertex 2*), maka penelitian yang dilakukan mendapatkan hasil yaitu 10 aktor kunci seperti pada gambar 9, serta graf yang divisualisasikan berdasarkan 10 aktor kunci pada gambar 10.

Tweet	Normalized
Temen gue ada udah miskin, ngatain temen gue yang turun demo dan ikut kawal putusan MK kemaren pengangguran. Bego. Emang gak semua orang kuliah punya otak sih.	teman saya sudah miskin hina teman saya turun demo kawal putus mahkamah konstitusi kemarin anggur bego memang tidak orang kuliah otak
@JhonSitorus_18 Jokowi "sembunyikan" anaknya, beda kelas dengan Wanda Hamidah yang bangga dengan anak-anaknya yang ikut demo #KawalPutusanMk https://t.co/QOQncuG71D #wandahamidah #fufufafa	sembunyi anak beda kelas bangga anak demo

Data yang sudah siap di analisis dapat dilihat pada tabel 1.

4.4. Tahap *Labeling*

Hasil *pre processing* dilabeling menggunakan *Inset Lexicon*. Penelitian ini menggunakan *Inset Lexicon* karena *Lexicon* dapat membantu memberi label secara otomatis dan algoritma SVM yang digunakan dalam penelitian ini membutuhkan data latih (*training data*) yang sudah berlabel untuk bisa membedakan antara kelas positif, negatif dan netral.

Tanpa adanya proses labeling , SVM tidak mengetahui mana data yang positif, negatif maupun netral, sehingga SVM tidak bisa belajar membentuk *hyperplane* pemisah antar kelas.

[illegible]

Gambar 11. Score sentimen

Kolom score pada gambar 11 digunakan untuk klasifikasi selanjutnya. Karena klasifikasi yang digunakan dalam penelitian ini membutuhkan input numerik.

4.5. TF-IDF

```
Jumlah data latih: 1012
Jumlah data uji: 254
Total data: 1266
```

Gambar 12. Jumlah *split* data

Data berlabel kemudian di *split* menjadi data latih dan data uji. Jumlah data terdapat pada gambar 12.

[illegible]

Gambar 13. *TF - IDF* data latih

Pembagian data menggunakan perbandingan 80 : 20. Data latih sebanyak 80% dari keseluruhan data seperti pada gambar 13.

[illegible]

Gambar 14. *TF-IDF* data uji

Data latih sebanyak 20% dari keseluruhan data seperti pada gambar 14.

Sebelum melakukan klasifikasi, penelitian ini menggunakan TF-IDF untuk mengubah teks menjadi numerik. Hasil pembobotan TF-IDF dalam penelitian ini terdapat pada gambar 13 dan 14.

4.6. Klasifikasi Sentimen

```
Confusion Matrix:
[[116   7   7]
 [ 39  21   7]
 [ 11  12  34]]
```

Gambar 15. *Confusion matrix* (kernel linear)

Confusion matrix digunakan untuk mendapatkan hasil pengujian dari model SVM terdapat pada gambar 15. Kernel yang digunakan dalam data adalah kernel linear.

Perhitungan pengujian terdapat pada gambar 16 yang memperlihatkan bahwa kernel linear mendapatkan akurasi sebesar 0.67 atau 67%.

Classification Report:				
	precision	recall	f1-score	support
negative	0.70	0.89	0.78	130
neutral	0.53	0.31	0.39	67
positive	0.71	0.60	0.65	57
accuracy			0.67	254
macro avg	0.64	0.60	0.61	254
weighted avg	0.66	0.67	0.65	254

Gambar 16. Hasil pengujian (kernel linear)

```

===== Kernel: RBF =====
Accuracy: 0.6299212598425197
Recall: 0.6299212598425197
Precision: 0.6697434450664326
F1-Score: 0.5655646358444147

Classification Report:
              precision    recall  f1-score   support

   negative         0.60         0.95         0.74         130
    neutral         0.73         0.12         0.21          67
    positive         0.76         0.49         0.60          57

   accuracy                    0.63         254
  macro avg         0.70         0.52         0.51         254
 weighted avg         0.67         0.63         0.57         254

Confusion Matrix:
[[124  2  4]
 [ 54  8  5]
 [ 28  1 28]]

```

Gambar 17. Hasil pengujian kernel RBF

Kernel RBF pada gambar 17 mendapatkan akurasi sebesar 0.63 atau 63%.

```

===== Kernel: POLY =====
Accuracy: 0.5866141732283464
Recall: 0.5866141732283464
Precision: 0.5856529244067759
F1-Score: 0.49806869984553365

Classification Report:
              precision    recall  f1-score   support

   negative         0.57         0.98         0.72         130
    neutral         0.44         0.06         0.11          67
    positive         0.78         0.32         0.45          57

   accuracy                    0.59         254
  macro avg         0.60         0.45         0.43         254
 weighted avg         0.59         0.59         0.50         254

```

Gambar 18. Hasil pengujian kernel polynomial

Kernel RBF pada gambar 18 mendapatkan akurasi sebesar 0.59 atau 59%.

```

===== Kernel: SIGMOID =====
Accuracy: 0.6377952755905512
Recall: 0.6377952755905512
Precision: 0.6124715350297525
F1-Score: 0.5973931898763961

Classification Report:
              precision    recall  f1-score   support

   negative         0.65         0.90         0.75         130
    neutral         0.48         0.19         0.28          67
    positive         0.68         0.56         0.62          57

   accuracy                    0.64         254
  macro avg         0.60         0.55         0.55         254
 weighted avg         0.61         0.64         0.60         254

Confusion Matrix:
[[117  6  7]
 [ 46 13  8]
 [ 17  8 32]]

```

Gambar 19. Hasil pengujian (Kernel Sigmoid)

Kernel RBF pada gambar 17 mendapatkan akurasi sebesar 0.64 atau 64%.

Evaluasi akurasi klasifikasi sentimen menggunakan SVM dengan empat jenis kernel menunjukkan bahwa kernel linear memberikan hasil terbaik dengan akurasi 67%, diikuti oleh kernel sigmoid (64%), RBF (63%), dan polinomial (59%).

Hasil ini menunjukkan bahwa dalam kasus analisis sentimen berbasis teks, pemetaan fitur dalam ruang vektor cenderung linier, sehingga kernel linear lebih efektif dibandingkan kernel non-linear seperti RBF dan polynomial.

Persentase hasil pengujian klasifikasi sentimen dari setiap kernel dihitung menggunakan rumus berikut

$$\text{Sentimen} = \frac{\text{Total sentimen}}{\text{Total data uji}} \times 100 \quad (11)$$

Kernel linear, RBF, polinomial, dan sigmoid menghasilkan jumlah sentimen yang sama. Maka perhitungan persentase didapatkan sebesar

$$\begin{aligned} \text{Negatif} &= \frac{\text{Total sentimen}}{\text{Total data uji}} \times 100 \\ &= \frac{130}{254} \times 100 \\ &= 51,1811024\% \quad (12) \end{aligned}$$

$$\begin{aligned} \text{Netral} &= \frac{\text{Total sentimen}}{\text{Total data uji}} \times 100 \\ &= \frac{67}{254} \times 100 \\ &= 26,377953\% \quad (12) \end{aligned}$$

$$\begin{aligned} \text{Positif} &= \frac{\text{Total sentimen}}{\text{Total data uji}} \times 100 \\ &= \frac{57}{254} \times 100 \\ &= 22,440945\% \quad (13) \end{aligned}$$

5. KESIMPULAN DAN SARAN

Media sosial X merupakan platform yang cukup efektif untuk penyebaran informasi. Dalam penelitian ini yang menganalisis SNA dan sentimen, didapatkan dua hasil dalam penelitian ini. Pertama, aktor yang berperan penting dalam penyebaran informasi terkait Peringatan Darurat RUU Pilkada di X adalah dessertmeys, dimm_sky, zyze9, hasbil_lbs, lalalabloem, dsperdana, ddazling, mamtalkatiri, tempodotco, divhumas_polri yang diurutkan dari teratas menurut hasil *Degree Centrality*. Kedua, dalam hal analisis sentimen, hasil yang didapatkan penelitian ini berupa akurasi penggunaan klasifikasi SVM dan tiga kategori sentimen yaitu positif, negatif, dan netral. Menurut hasil penelitian, akurasi yang didapatkan penelitian ini dengan penggunaan kernel linear, sigmoid, RBF, dan polinomial yaitu 67%, 64%, 63%, dan 59% secara berurutan dan dari total data uji, sebanyak 51,18% pengguna merespon dengan sentimen negatif, 26,38% memberi tanggapan netral dan 22,44% bertanggapan positif. Saran berdasarkan penelitian yang dilakukan ini adalah mencari cara *labeling* yang tepat untuk data sehingga hasilnya lebih

akurat serta menemukan algoritma yang tepat sesuai dengan data yang dihasilkan.

DAFTAR PUSTAKA

- [1] Rahmadhany, A., Aldila Safitri, A., & Irwansyah, I. (2021). Fenomena Penyebaran Hoax dan Hate Speech pada Media Sosial. *Jurnal Teknologi Dan Sistem Informasi Bisnis*, 3(1), 30–43. <https://doi.org/10.47233/jteksis.v3i1.182>
- [2] Yusuf, F., Rahman, H., Rahmi, S., Lismayani, A., & Guru Sekolah Dasar Universitas Negeri Makassar, P. (2023). Pemanfaatan Media Sosial sebagai Sarana Komunikasi, Informasi, dan Dokumentasi: Pendidikan di Majelis Taklim Annur Sejahtera. *Jurnal Hasil-Hasil Pengabdian Dan Pemberdayaan Masyarakat*, 2. <https://journal.unm.ac.id/index.php/JHP2M>
- [3] Ginting, J. A., Manongga, D., & Sembiring, I. (n.d.). *The Spread Path of Hoax News in Social Media (Facebook) using Social Network Analysis (SNA)*.
- [4] Sembiring, I., Iriani, A., Veronika Br Ginting, J., & Amanda Ginting, J. (n.d.). A Novel Approach to Network Forensic Analysis: Combining Packet Capture Data and Social Network Analysis. In *IJACSA) International Journal of Advanced Computer Science and Applications* (Vol. 14, Issue 3). www.ijacsa.thesai.org
- [5] Fitriyah, N., Warsito, B., Asih, D., & Maruddani, I. (2020). ANALISIS SENTIMEN GOJEK PADA MEDIA SOSIAL TWITTER DENGAN KLASIFIKASI SUPPORT VECTOR MACHINE (SVM). *JURNAL GAUSSIAN*, 9(3), 376–390. <https://ejournal3.undip.ac.id/index.php/gaussian/>
- [6] Putri, A. J., Syafira, A. S., Purbaya, M. E., & Purnomo, D. (2022). Analisis Sentimen E-Commerce Lazada pada Jejaring Sosial Twitter Menggunakan Algoritma Support Vector Machine. *Jurnal TRINISTIK: Jurnal Teknik Industri, Bisnis Digital, Dan Teknik Logistik*, 1(1), 16–21. <https://doi.org/10.20895/trinistik.v1i1.447>
- [7] Gifari, O. I., Adha, M., Rifky Hendrawan, I., Freddy, F., & Durrand, S. (2022). Analisis Sentimen Review Film Menggunakan TF-IDF dan Support Vector Machine. *JIFOTECH (JOURNAL OF INFORMATION TECHNOLOGY)*, 2(1).
- [8] Br Sinulingga, J. E., & Sitorus, H. C. K. (2024). Analisis Sentimen Masyarakat terhadap Film Horor Indonesia Menggunakan Metode SVM dan TF-IDF. *Jurnal Manajemen Informatika (JAMIKA)*, 14(1), 42–53. <https://doi.org/10.34010/jamika.v14i1.11946>
- [9] Fikri, M. I., Sabrila, T. S., Azhar, Y., & Malang, U. M. (n.d.). *Perbandingan Metode Naïve Bayes dan Support Vector Machine pada Analisis Sentimen Twitter*.
- [10] Fernanda, M., & Fathoni, N. (2024). Perbandingan Performa Labeling Lexicon InSet dan VADER pada Analisa Sentimen Rohingya di Aplikasi X dengan SVM. *Jurnal Informatika Dan Sains Teknologi*, 1(3), 62–76. <https://doi.org/10.62951/modem.v1i3.112>
- [11] Prastyo, P. H., Ardiyanto, I., & Hidayat, R. (2020, October 26). Indonesian Sentiment Analysis: An Experimental Study of Four Kernel Functions on SVM Algorithm with TF-IDF. *2020 International Conference on Data Analytics for Business and Industry: Way Towards a Sustainable Economy, ICDABI 2020*. <https://doi.org/10.1109/ICDABI51230.2020.9325685>
- [12] Rahayu, I.P., Fauzi, A., & Indra, J. (2022). Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Naive Bayes dan Support Vector Machine. *Jurnal Sistem Komputer dan Informatika (JSON)*, 4(2), 293–301. <https://doi.org/10.30865/json.v4r2.5381>
- [13] Giharis, N., & Herdian, C (2023). *Sentiment Analysis of The Community's Response Towards The Magical Ambassador Brand Lisa Blackpink on Twitter Social Media using The Naive Bayes Method*. *Social Science Research Network (SSRN)*