

Analisis Sentimen Publik Terhadap Penegakan Hukum dalam Penangkapan Ikan Ilegal oleh Kapal Vietnam di Twitter Menggunakan Metode Graph-based Reasoning and Analysis for Text (GRAFT)

Rama Ardiansyah¹, Nur Inayah², Muhaza Liebenlito³

ramaardiansyah.mtk@gmail.com, nur.inayah@uinjkt.ac.id, muhazaliebenlito@uinjkt.ac.id

^{1,2,3}Universitas Islam Negeri Syarif Hidayatullah Jakarta, Indonesia

Informasi Artikel	Abstrak
Diterima : 15 Jan 2025 Direview : 25 Feb 2025 Disetujui : 15 Apr 2025	Penangkapan ikan ilegal oleh kapal asing, khususnya kapal Vietnam, merupakan isu utama di Indonesia yang merugikan ekonomi dan mengancam keberlanjutan sumber daya laut. Penegakan hukum terhadap praktik ini sangat penting untuk menjaga sumber daya alam dan kestabilan ekonomi. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap penegakan hukum dalam penangkapan ikan ilegal menggunakan metode Graph-based Reasoning and Analysis for Text (GRAFT). Data yang digunakan diambil dari Twitter dengan teknik crawling, mencakup 1.228 tweet. Penelitian ini menerapkan Graph Convolutional Network (GCN) untuk klasifikasi sentimen dan evaluasi menggunakan K-Fold Cross Validation. Hasil penelitian menunjukkan akurasi model 96.61%, dengan Precision 96.60%, Recall 96.61%, dan F1-Score 96.60%. Meskipun demikian, keterbatasan penelitian ini terletak pada data yang hanya bersumber dari satu platform. Penelitian selanjutnya dapat melibatkan data dari berbagai platform media sosial dan pendekatan multibahasa untuk hasil yang lebih komprehensif.
Kata Kunci	
Penangkapan ikan ilegal, Kapal Vietnam, Twitter, Sentimen publik, Metode GRAFT	
Keywords	Abstract
Illegal fishing, Vietnamese vessels, Twitter, Public sentiment, GRAFT method	<i>Illegal fishing by foreign vessels, particularly those from Vietnam, is a major issue in Indonesia, harming the economy and threatening marine resource sustainability. Law enforcement against this activity is crucial for protecting natural resources and maintaining economic stability. This study aims to analyze public sentiment towards law enforcement in illegal fishing using the Graph-based Reasoning and Analysis for Text (GRAFT) method. Data was collected from Twitter using crawling techniques, covering 1,228 tweets. The study applied Graph Convolutional Network (GCN) for sentiment classification, with evaluation using K-Fold Cross Validation. The results show that the model achieved an accuracy of 96.61%, with Precision 96.60%, Recall 96.61%, and F1-Score 96.60%. However, the limitation of this research lies in the data sourced from a single social media platform. Future research could involve data from various social media platforms and a multilingual approach for more comprehensive results.</i>

A. Pendahuluan

Penangkapan ikan ilegal merupakan permasalahan yang berdampak luas terhadap sumber daya laut, ekonomi, dan keamanan negara [1]. Di Indonesia, kapal-kapal asing, khususnya dari Vietnam, sering kali terlibat dalam aktivitas ini. Berbagai metode analisis sentimen telah digunakan untuk memahami opini publik, termasuk Naïve Bayes, Support Vector Machine (SVM), Random Forest, dan Deep Learning [6]. Namun, metode-metode tersebut memiliki keterbatasan dalam menangkap hubungan semantik antar kata dalam teks.

Metode GRAFT (Graph-based Reasoning and Analysis for Text) menawarkan pendekatan berbasis graf yang lebih fleksibel dalam menangani data tekstual dalam jumlah besar, seperti tweet [7]. Dengan teknik ini, hubungan antar kata direpresentasikan dalam bentuk node dan edge, memungkinkan pemetaan konteks sentimen secara lebih akurat. Oleh karena itu, penelitian ini mengusulkan penggunaan GRAFT untuk analisis sentimen terkait penangkapan ikan ilegal oleh kapal Vietnam.

Pada penelitian ini, GRAFT diterapkan bersama dengan Graph Convolutional Network (GCN) untuk klasifikasi sentimen. Proses ini mencakup preprocessing data, pembentukan graf, pelatihan model, dan evaluasi performa dengan Precision, Recall, F1-score, serta Cross Validation [8]. Hasil penelitian ini diharapkan dapat memberikan wawasan bagi pemerintah dalam merumuskan kebijakan terkait penanganan Illegal, Unreported, and Unregulated (IUU) Fishing.

B. Metode Penelitian

1. Pengumpulan Data

Data diperoleh dari Twitter melalui crawling, fokus pada tweet tentang penangkapan ikan ilegal oleh kapal Vietnam. Pengumpulan berlangsung dari 27 November 2014 hingga 13 Maret 2024, menggunakan kata kunci seperti #PenangkapanIkanIllegalVietnam. Dikumpulkan 1.228 tweet dalam Bahasa Indonesia dengan engagement tinggi. Filtering berulang dilakukan untuk memastikan data relevan dan akurat.

Tabel 1 Data Hasil Crawling

Atribut	Deskripsi	Jumlah data
created_at	Tanggal dan waktu tweet diposting	1.228
full_text	Isi lengkap dari tweet	1.228
lang	Bahasa yang digunakan dalam tweet	1.228
username	Nama pengguna Twiter yang memposting tweet	1.228
favorite_count	Jumlah “likes”	1.228
retweet_count	Jumlah retweet	1.228
location	Lokasi pengguna	928
tweet_url	URL dari tweet yang diposting	1.228

Dari tabel 1 diperoleh data yang terdiri dari 1.228 tweet dari berbagai pengguna dengan atribut penting, termasuk `full_text` yang menunjukkan isi lengkap dari tweet, `created_at` yang merujuk pada tanggal dan waktu tweet diposting, dan `username` yang merupakan nama pengguna Twitter. Hasil dari data yang telah difilter kemudian dianalisis untuk memahami tren terkait penangkapan ikan ilegal oleh kapal Vietnam.

2. Desain Penelitian

Metode GRAFT (Graph-based Reasoning and Analysis for Text) dipilih karena memiliki kemampuan unggul dalam menganalisis data teks dengan tingkat kedalaman melalui pemanfaatan representasi grafis. Dengan GRAFT, hubungan antar kata dan frasa dalam teks, seperti tweet, dapat diidentifikasi dan dipetakan secara visual dalam bentuk graf yang kompleks. Graf dapat membantu pemahaman yang lebih mendalam tentang hubungan antar unit teks. Dengan menggunakan struktur graf, keterkaitan antara elemen-elemen teks dapat diidentifikasi dan dianalisis secara lebih terstruktur dan jelas[9]. Hal ini memungkinkan peneliti untuk mengungkap pola-pola sentimen yang tersembunyi, serta untuk memahami konteks dan makna yang lebih dalam dari setiap interaksi antar kata dalam teks.

3. Cleaning Tweet

Prosedur dimulai dengan preprocessing data, Tujuan dari pre-processing adalah untuk mengubah data mentah atau data yang tidak teratur menjadi lebih bersih, mengelompokkan data, dan mempersiapkannya agar lebih mudah dianalisis dan diproses oleh suatu algoritma atau metode perhitungan[10]. Pada proses preprocessing meliputi normalisasi, tokenizing, penghapusan stopwords, dan stemming. Tahapan ini dimulai dengan normalisasi, yaitu mengubah teks menjadi format yang konsisten, seperti mengonversi semua huruf menjadi huruf kecil, menghapus emoji dan simbol, serta menghapus kata-kata yang tidak relevan. Langkah berikutnya adalah tokenisasi, untuk memecah teks menjadi unit-unit kecil berupa kata atau frasa. Penghapusan stopwords digunakan untuk menghapus kata-kata umum seperti "dan", "atau", "yang", dan kata-kata umum lainnya yang dianggap tidak memiliki nilai informatif. Terakhir, proses stemming digunakan untuk mengubah kata-kata ke bentuk dasarnya guna mengurangi variasi kata dengan makna serupa.

4. Pelabelan

Penelitian ini menggunakan pelabelan otomatis berbasis leksikal (*unsupervised labeling*) dengan pendekatan *SentiWordNet*. Metode ini memungkinkan penentuan label sentimen tanpa perlu anotasi manual, sehingga cocok digunakan dalam dataset yang belum memiliki label sentimen. Dalam pendekatan ini, setiap kata dalam teks dievaluasi berdasarkan nilai sentimen yang diberikan oleh *SentiWordNet*, yang mencakup nilai positif dan negatif.

Proses pelabelan dimulai dengan prapemrosesan teks, yang mencakup tokenisasi, penghapusan stopwords, serta penghapusan tanda baca. Setelah teks diproses, setiap kata dalam teks akan dicocokkan dengan daftar kata dalam *SentiWordNet*. Jika kata tersebut memiliki entri dalam *SentiWordNet*, maka sistem akan mengambil nilai positif dan negatif dari kata tersebut. Skor sentimen dari suatu

kata dihitung berdasarkan selisih antara skor positif dan negatifnya, yang dirumuskan sebagai berikut:

$$S(w) = \frac{\sum(pos_score(w) - neg_score(w))}{N} \quad (1)$$

Di mana:

- $S(w)$ adalah skor sentimen untuk kata w
- $pos_score(w)$ adalah skor positif dari kata w dalam *SentiWordNet*
- $neg_score(w)$ adalah skor negatif dari kata w dalam *SentiWordNet*
- N adalah jumlah sinonim kata w dalam *SentiWordNet*

Setelah skor sentimen untuk setiap kata dihitung, nilai sentimen dari seluruh kata dalam teks dijumlahkan untuk mendapatkan skor sentimen total dari teks, yang diformulasikan sebagai:

$$S(T) = \sum_{i=1}^M S(w_i) \quad (2)$$

Di mana:

- $S(T)$ adalah skor sentimen total dari teks T
- M adalah jumlah kata dalam teks T
- $S(w_i)$ adalah skor sentimen dari kata ke- i dalam teks

Berdasarkan skor sentimen total $S(T)$ yang diperoleh, label sentimen ditentukan menggunakan aturan berikut:

$$Label(T) = \begin{cases} Positive, & \text{jika } S(T) \geq 0 \\ Negative, & \text{jika } S(T) < 0 \end{cases} \quad (3)$$

Dengan demikian, jika skor total lebih besar atau sama dengan nol, teks dikategorikan sebagai positif, sedangkan jika skor total kurang dari nol, teks dikategorikan sebagai negatif.

Metode ini memungkinkan klasifikasi sentimen dilakukan secara otomatis dan konsisten berdasarkan sumber leksikal. Namun, pendekatan berbasis leksikon memiliki keterbatasan dalam menangani konteks kalimat, sarkasme, serta pengaruh kata negasi yang dapat membalikkan makna suatu teks. Oleh karena itu, dalam penelitian ini, pendekatan ini dikombinasikan dengan representasi berbasis Graph Convolutional Networks (GCN) untuk menangkap hubungan antar kata dan meningkatkan akurasi klasifikasi sentimen.

5. Graph Convolution Network (GCN)

GCN dikembangkan dari *graph neural network* (GNN). Definisi GNN, kita secara formal mendefinisikan himpunan *graph* $G = (g_1, g_2, \dots, g_n)$. Untuk setiap grafik $g = (V, E)$, di mana V adalah himpunan *graph*, E adalah himpunan sisi pada *graph* (edge) tersebut. Setiap sembarang node $v \in V$, terdapat sisi yang terhubung dengan dirinya sendiri, yaitu $(v, v) \in E$. Untuk setiap *graph* terdapat matriks $X \in \mathbb{R}^{i \times j}$ yang merepresentasikan himpunan ciri awal setiap node pada graf tersebut, dimana i mewakili jumlah node dalam *graph* dan j mewakili dimensi fitur setiap node. Kami akan membahas bagaimana mendapatkan keadaan awal pada tahap selanjutnya.

Setiap *graph* terdapat matriks ketetanggaan A dan matriks derajat D yang bersesuaian. Maka GCN dapat di definisikan rumus lapisan pertama adalah :

$$H^{(1)} = \text{ReLU}(\tilde{A}XW_0) \quad (4)$$

Rumus diatas, $W_0 \in \mathbb{R}^{i \times j}$ mewakili matriks bobot. $\text{ReLU}(\cdot)$ mewakili fungsi aktivasi, didefinisikan sebagai $\text{ReLU}(x) = \max(0, x)$. Iterasi jumlah lapisan GCN bisa mendapatkan fitur yang lebih dalam. Dengan menggunakan rumus pembaharuan status tersembunyi GCN diperoleh sebagai berikut :

$$H^{(t+1)} = \text{ReLU}(\tilde{A}H^{(t)}W_0) \quad (5)$$

Dimana t menyatakan jumlah iterasi. Secara khusus $H(0) = X$.

6. Pendekatan Graf

Graf terdiri dari simpul-simpul (nodes) yang merepresentasikan berbagai entitas. Sementara itu, tepi (edges) pada graf menggambarkan hubungan atau interaksi antara entitas-entitas tersebut, menciptakan jaringan yang kompleks dan terstruktur[12]. Model GRAFT diterapkan untuk mengidentifikasi pola sentimen dalam jaringan kata tersebut, yang kemudian dikategorikan sebagai sentimen positif dan negatif. Interpretasi hasil dilakukan dengan mengaitkan sentimen yang muncul dengan peristiwa atau kebijakan yang relevan selama periode pengumpulan data. Berikut ini langkah-langkah proses model GRAFT:

a. Pembuatan Graf Teks

Langkah ini bertujuan untuk merepresentasikan hubungan antar kata dalam bentuk graf. Setiap kata dalam teks dijadikan simpul (*node*), sedangkan hubungan antar kata yang muncul secara berurutan dalam teks diwakili sebagai tepi (*edge*). Dengan menggunakan pustaka *NetworkX*, graf dibangun dengan menambahkan setiap kata ke dalam simpul dan menghubungkannya dengan kata sebelumnya sebagai tepi.

b. Analisis dan Visualisasi Graf

Graf yang telah dibangun divisualisasikan untuk menggambarkan hubungan antar kata secara struktural. Menggunakan layout Kamada-Kawai, node (kata) diatur dalam posisi yang rapi dan saling terkait dengan edge (hubungan antar kata). Visualisasi dilakukan menggunakan *matplotlib* dan *pyvis*, yang memungkinkan pengamatan pola hubungan antar kata secara interaktif.

c. Pembagian Data

Sebelum dilakukan pelatihan model, dataset dibagi menjadi dua bagian, yaitu data latih dan data uji. Pembagian dilakukan dengan menggunakan rasio 80:20, di mana 80% data digunakan untuk melatih model, sedangkan 20% sisanya digunakan sebagai data uji untuk mengevaluasi kinerja model. Proses pembagian data ini dilakukan menggunakan *train_test_split* dari *scikit-learn*, dengan *random_state=42* untuk memastikan reproduktibilitas hasil. Pembagian data ini bertujuan untuk memastikan bahwa model diuji pada data yang tidak terlihat selama pelatihan, sehingga dapat mengukur kemampuannya dalam menggeneralisasi pada data yang baru.

d. Implementasi Model *Graph Convolutional Network* (GCN)

Graph Convolutional Network (GCN) adalah model deep learning berbasis graf yang digunakan untuk mengolah data dengan struktur graf. Dalam penelitian ini, GCN digunakan untuk analisis sentimen berbasis graf, di mana

kata-kata dalam teks direpresentasikan sebagai node, dan hubungan antar kata dalam kalimat yang sama direpresentasikan sebagai edge dalam graf. Dengan pendekatan ini, GCN dapat menangkap hubungan semantik antar kata yang lebih kompleks dibandingkan dengan model berbasis kata individu seperti bag-of-words atau TF-IDF.

1) Konversi Teks ke Graf

Langkah pertama dalam implementasi GCN adalah mengubah teks menjadi graf menggunakan NetworkX. Setiap teks dalam dataset diubah menjadi graf, di mana setiap kata diperlakukan sebagai node dan setiap pasangan kata yang berdekatan dalam kalimat dihubungkan dengan edge. Bobot pada setiap edge dihitung berdasarkan frekuensi kemunculan pasangan kata dalam teks.

2) Struktur Model GCN

Model GCN dalam penelitian ini terdiri dari tiga lapisan utama:

- a) Lapisan Input: Menerima vektor fitur awal untuk setiap node, yang direpresentasikan dalam bentuk one-hot encoding.
- b) Lapisan Konvolusi Berbasis Graf: Memanfaatkan operasi Graph Convolutional Layer untuk menyebarkan informasi antar node tetangga. Informasi dari node tetangga digabungkan untuk memperbarui representasi node saat propagasi.
- c) Lapisan Output: Menggunakan fungsi Softmax untuk menghasilkan prediksi kelas sentimen, baik positif maupun negatif.

Proses propagasi dalam model GCN mengikuti rumus: *Graph Convolutional Network* (GCN)

$$H^{(l+1)} = \sigma \left(\tilde{D}^{-\frac{1}{2}} \tilde{A} \tilde{D}^{-\frac{1}{2}} H^{(l)} W^{(l)} \right) \quad (6)$$

Keterangan :

- $H^{(l)}$ adalah representasi fitur dari node pada lapisan ke- l .
- $\tilde{A} = A + I$ adalah matriks adjacency dengan self-loop yang menambahkan koneksi pada setiap node dengan dirinya sendiri.
- $\tilde{D}$ adalah matriks derajat diagonal dari $\tilde{A}$, yang digunakan untuk normalisasi.
- $W^{(l)}$ adalah bobot pada lapisan ke- l , yang dipelajari selama pelatihan.
- σ adalah fungsi aktivasi ReLU yang digunakan untuk memastikan propagasi informasi tetap positif.

Lapisan konvolusi ini bekerja dengan menyebarkan informasi dari node tetangga dan memperbarui representasi node berdasarkan interaksi antar node dalam graf.

3) Fungsi Loss dan Optimasi

Model ini menggunakan Cross-Entropy Loss untuk mengukur seberapa jauh prediksi model dari label asli. Fungsi loss ini dirumuskan sebagai:

$$\mathcal{L} = - \sum_i y_i \log(\hat{y}_i) \quad (7)$$

di mana :

- y_i adalah label sentimen asli dari node ke- i .
- $\hat{y}_i$ adalah probabilitas prediksi model terhadap node tersebut.

Untuk mengoptimalkan model, digunakan Adam Optimizer, yang dipilih karena kemampuannya dalam menangani data berbasis graf yang sparsitasnya tinggi. Adam menggabungkan keuntungan dari Momentum Optimization dan RMSProp, sehingga lebih stabil dalam pelatihan. Pembaruan parameter dilakukan menggunakan rumus:

$$\theta_t = \theta_{t-1} - \alpha \frac{\hat{m}_t}{\sqrt{\hat{v}_t} + \epsilon} \quad (8)$$

Keterangan:

- α adalah learning rate.
- $\hat{m}_t$ adalah estimasi pertama dari gradien (momentum).
- $\hat{v}_t$ adalah estimasi kedua dari gradien (kecepatan perubahan gradien).
- ϵ adalah nilai kecil untuk mencegah pembagian dengan nol.

e. Implementasi Cross-Validation dengan K-Fold

Untuk memastikan bahwa model tidak hanya bekerja dengan baik pada satu subset data tertentu, dilakukan K-Fold Cross-Validation. Dalam teknik ini, dataset dibagi menjadi K bagian (folds), dan model dilatih pada K-1 bagian sambil diuji pada 1 bagian yang tersisa. Proses ini dilakukan sebanyak K kali, sehingga setiap bagian digunakan sebagai data uji sekali. Teknik ini membantu menghindari bias dan memastikan model dapat menggeneralisasi dengan baik pada data yang belum pernah dilihat sebelumnya.

Dalam penelitian ini, dilakukan cross-validation dengan nilai K yang bervariasi (5, 10, 15, 20, 25, 30) untuk mengevaluasi kinerja model dalam berbagai kondisi pembagian data. Hasil evaluasi dari setiap fold dihitung rata-ratanya untuk memperoleh metrik kinerja akhir.

Proses cross-validation dilakukan dengan membagi data menjadi K bagian (folds) yang berbeda. Model kemudian dilatih menggunakan K-1 bagian dari dataset, sementara satu bagian sisanya digunakan untuk pengujian. Proses ini diulang hingga setiap bagian digunakan sebagai data uji. Setelah itu, hasil evaluasi dari semua fold dihitung rata-ratanya. Metrik evaluasi yang digunakan dalam penelitian ini meliputi akurasi, precision, recall, dan F1-score.

f. Evaluasi dan Visualisasi Hasil Menggunakan data test

Setelah model dilatih dan divalidasi dengan K-Fold Cross-Validation, evaluasi akhir dilakukan menggunakan data test yang telah dipisahkan sebelumnya untuk mengukur performa model dalam mengklasifikasikan sentimen pada data baru. Metrik evaluasi yang digunakan meliputi akurasi, precision, recall, dan F1-score, serta Confusion Matrix yang divisualisasikan dengan heatmap untuk mempermudah interpretasi. Selain evaluasi numerik, analisis sentimen juga divisualisasikan dalam bentuk Word Cloud untuk menampilkan kata-kata yang dominan dalam sentimen positif dan negatif. Proses evaluasi ini mencakup pengujian model pada data test, perhitungan serta visualisasi Confusion Matrix, dan pembuatan Word Cloud. Evaluasi ini memastikan model mampu menggeneralisasi dengan baik dan memahami pola hubungan antar kata dalam teks.

Berikut adalah penjelasan dan rumus perhitungan untuk setiap metrik evaluasi:

$$Akurasi = \frac{TP + TN}{TP + TN + FP + FN} \quad (9)$$

$$Precision = \frac{TP}{TP + FP} \quad (10)$$

$$Recall = \frac{TP}{TP + FN} \quad (11)$$

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (12)$$

Keterangan :

- TP = Banyaknya data positif yang benar
- TN = Banyaknya data negatif yang benar
- FN = Banyaknya data negatif yang salah
- FP = Banyaknya data positif yang salah

C. Hasil Dan Pembahasan

1. Preprocessing Data

Pada tahap ini, teks yang digunakan dalam analisis dipersiapkan untuk diolah. Proses preprocessing ini melibatkan beberapa langkah penting seperti konversi teks menjadi huruf kecil, penghapusan tanda baca, serta penghapusan stopwords. Penghapusan stopwords sangat penting untuk mengurangi kata-kata yang tidak memberikan informasi signifikan dalam analisis sentimen. Berikut adalah langkah-langkah yang dilakukan dalam preprocessing data:

- Pengubahan teks menjadi huruf kecil untuk memastikan keseragaman dalam pemrosesan teks.
- Penghapusan tanda baca dengan menggunakan fungsi translate untuk memudahkan pemisahan kata-kata.
- Tokenisasi teks dengan memecah kalimat menjadi kata-kata.
- Penghapusan stopwords untuk menghilangkan kata-kata umum yang tidak membawa makna, seperti “dan”, “atau”, yang tidak membantu dalam analisis sentimen.

Tabel dibawah ini bertujuan untuk menunjukkan bagaimana data mentah yang awalnya berupa teks bebas (tweet) diubah menjadi format yang lebih terstruktur dan siap untuk analisis lebih lanjut, khususnya untuk analisis sentimen. Memberikan wawasan kepada pembaca tentang bagaimana teks diproses untuk mengurangi noise dan hanya menyisakan informasi yang relevan, yang akan digunakan untuk model analisis sentimen.

Tabel 2 Hasil Preprocessing Data

Teks Asli	Hasil Preprocessing
Amankan Kapal Ilegal Fishing di Laut Sulawesi Polri melalui Polda Sulut berhasil mengamankan kapal ilegal fishing di perairan Indonesia. Kapal asal Filipina yang bernama Queen Davie tersebut melakukan penangkapan ikan tanpa dokumen perizinan dari pihak terkait. https://t.co/CM3wZahRRo	['aman', 'kapal', 'ilegal', 'fishing', 'laut', 'sulawesi', 'polri', 'lalu', 'polda', 'sulut', 'hasil', 'aman', 'kapal', 'ilegal', 'fishing', 'air', 'indonesia', 'kapal', 'asal', 'filipina', 'nama', 'queen', 'davie', 'sebut', 'laku', 'tangkap', 'ikan', 'dokumen', 'izin', 'pihak', 'kait']
@bambangmythic Dengan NCC Indonesia akan dapat lebih efektif dalam melawan kegiatan ilegal seperti penangkapan ikan ilegal dan penyelundupan narkoba.	['ncc', 'indonesia', 'lebih', 'efektif', 'lawan', 'giat', 'ilegal', 'tangkap', 'ikan', 'ilegal', 'selundup', 'narkoba']
Penangkapan Ikan Ilegal Ancam Ekosistem Laut Aceh https://t.co/5ulXSoHtRs	['tangkap', 'ikan', 'ilegal', 'ancam', 'ekosistem', 'laut', 'aceh']
Ngelawak.... Yg berkuasa 9thn Partai pengusung dan kadernya jadi Presiden kok.... Menteri disebut yg bukan dari kader partai manapun https://t.co/DzcT4zWPdv Calon Presiden (capres) nomor urut 3 penanganan penangkapan ikan secara ilegal atau illegal fishing.	['ngelawak', 'kuasa', 'partai', 'usung', 'kader', 'jadi', 'presiden', 'kok', 'menteri', 'sebut', 'bukan', 'kader', 'partai', 'mana', 'calon', 'presiden', 'capres', 'nomor', 'urut', 'tangan', 'tangkap', 'ikan', 'ilegal', 'illegal', 'fishing']

2. Pelabelan Sentimen

Pelabelan sentimen dalam penelitian ini dilakukan untuk mengklasifikasikan opini yang terkandung dalam teks tweet terkait penangkapan ikan ilegal oleh kapal asing, khususnya kapal Vietnam. Proses pelabelan ini dilakukan dengan menggunakan metode berbasis skor sentimen yang dihitung dengan menggunakan SentiWordNet, yang menyediakan skor positif, negatif, dan netral untuk setiap kata dalam teks.

Berdasarkan hasil pelabelan, dataset yang digunakan dalam penelitian ini memiliki distribusi label yang tidak seimbang, di mana terdapat 906 tweet positif dan 322 tweet negatif, dengan rasio positif:negatif sebesar 74%:26%.

Tabel berikut menunjukkan jumlah tweet dengan masing-masing label sentimen:

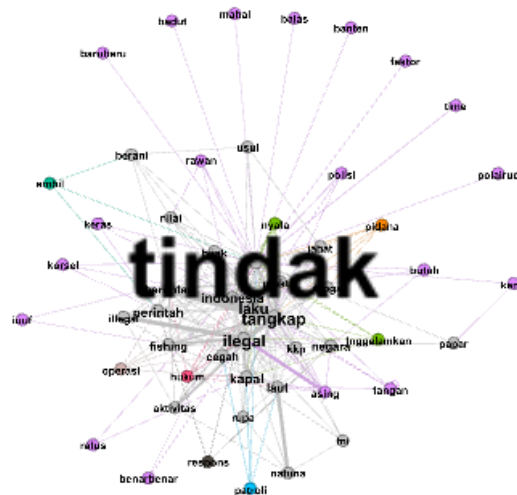
Tabel 3 Hasil Pelabelan

SENTIMEN LABEL	COUNT
POSITIVE	906
NEGATIVE	322

Pada gambar 2 menunjukkan kata-kata yang paling sering muncul dalam teks yang dikategorikan sebagai sentimen positif. Beberapa poin yang perlu diperhatikan:

- Kata Kunci Utama: Kata-kata yang dominan dalam word cloud ini seperti "tangkap", "ikan", "illegal", dan "Indonesia" juga muncul dengan frekuensi yang tinggi, mirip dengan word cloud untuk sentimen negatif.
- Interpretasi: Meskipun kata-kata yang muncul serupa, konteks penggunaannya kemungkinan berbeda. Dalam konteks positif, kata-kata seperti "tangkap" dan "illegal" mungkin digunakan untuk memuji tindakan pihak berwenang dalam menangkap kapal-kapal ilegal, serta **penegakan hukum yang tegas** oleh otoritas seperti pemerintah Indonesia.
- Asosiasi Kontekstual: Kemunculan kata "Indonesia" dalam word cloud positif menunjukkan bahwa sentimen positif mungkin juga terkait dengan kebanggaan nasional atau dukungan terhadap kebijakan pemerintah Indonesia dalam mengatasi masalah ini.

4. Visualisasi Klasifikasi Sentimen



Gambar 3 Klasifikasi Sentimen Positif Dengan Model GCN

Pada gambar 3, setiap *node* atau titik mewakili kata yang muncul dalam tweet, sementara *edge* atau garis penghubung antara node menunjukkan hubungan atau asosiasi antara kata-kata tersebut berdasarkan kemunculan mereka dalam konteks yang sama. Kata "tindak" tampak sebagai node terbesar, yang menunjukkan bahwa kata ini adalah yang paling sering muncul atau paling dominan dalam kumpulan data tweet yang dianalisis. Garis-garis penghubung atau edge menunjukkan hubungan antar kata. Misalnya, kata "tindak" terhubung dengan kata-kata lain seperti "kapal", "illegal", dan "tangkap", yang mengindikasikan bahwa tweet yang dianalisis sering membahas tindakan terhadap kapal yang melakukan penangkapan ikan ilegal. Dengan memanfaatkan hubungan antar kata yang ditunjukkan dalam visualisasi di atas, GCN dapat mengidentifikasi pola-pola dan konteks yang cenderung menunjukkan sentimen positif terhadap tindakan yang diambil terhadap pelaku penangkapan ikan ilegal. Kata "tindak" sering muncul bersamaan dengan kata-kata yang mengindikasikan aksi tegas atau keberhasilan



Gambar 4 Klasifikasi Sentimen Negatif Dengan Model GCN

Pada gambar 4 kata "illegal" menjadi node utama yang paling menonjol dan dihubungkan dengan berbagai kata lainnya. Ini menunjukkan bahwa kata "illegal" sering muncul dan memiliki bobot yang besar dalam analisis sentimen ini. Garis-garis yang menghubungkan "illegal" dengan kata-kata lain mengindikasikan hubungan atau keterkaitan di antara kata-kata tersebut. Kata-kata seperti "fishing," "praktik," "tindakan," dan "unregulated" sering muncul bersama dengan "illegal," mencerminkan topik-topik negatif yang sering dibahas dalam tweet. Ketika kata "illegal" sering muncul bersama dengan kata-kata yang menunjukkan aktivitas yang melanggar hukum atau praktik yang merugikan, GCN dapat mengidentifikasi ini sebagai indikasi sentimen negatif.

5. Hasil K-Fold Cross Validation

Untuk memastikan kinerja model yang stabil dan handal, dilakukan evaluasi menggunakan K-Fold Cross Validation. Metode ini membagi dataset menjadi K bagian (folds) yang hampir sama besar. Model diuji dengan menggunakan satu fold sebagai data uji dan K-1 fold lainnya sebagai data latih. Proses ini diulang untuk setiap fold, memastikan bahwa setiap bagian data digunakan baik untuk pelatihan maupun pengujian, sehingga dapat memberikan gambaran yang lebih menyeluruh tentang kinerja model pada berbagai subset data.

Dalam penelitian ini, rentang nilai K yang digunakan adalah 5 hingga 30, dengan tujuan untuk memeriksa pengaruh jumlah fold terhadap performa model. Setiap nilai K diuji untuk melihat stabilitas dan akurasi model pada data yang berbeda-beda. Hasil evaluasi menunjukkan bahwa performa model cenderung meningkat seiring dengan bertambahnya nilai K, dengan K=30 memberikan hasil akurasi tertinggi.

Berikut adalah hasil evaluasi per fold untuk Akurasi, Precision, Recall, dan F1-Score pada setiap nilai K yang diuji:

Tabel 4 Hasil Cross Validation

K	Akurasi	Precision	Recall	F1-Score
5	0.6859	0.5980	0.6859	0.5605
10	0.7083	0.7647	0.7083	0.6075
15	0.8268	0.8557	0.8268	0.8028
20	0.9016	0.9048	0.9016	0.8982
25	0.9364	0.9376	0.9364	0.9355
30	0.9556	0.9566	0.9556	0.9551

Tabel 4 menunjukkan bahwa model GCN mengalami peningkatan kinerja seiring dengan meningkatnya nilai K dalam cross-validation. Untuk K=30, model mencapai Akurasi sebesar 95.56%, dengan F1-Score yang sangat baik yaitu 95.51%. Hasil ini menunjukkan bahwa model mampu mengklasifikasikan sentimen dengan akurat dan konsisten pada data yang bervariasi.

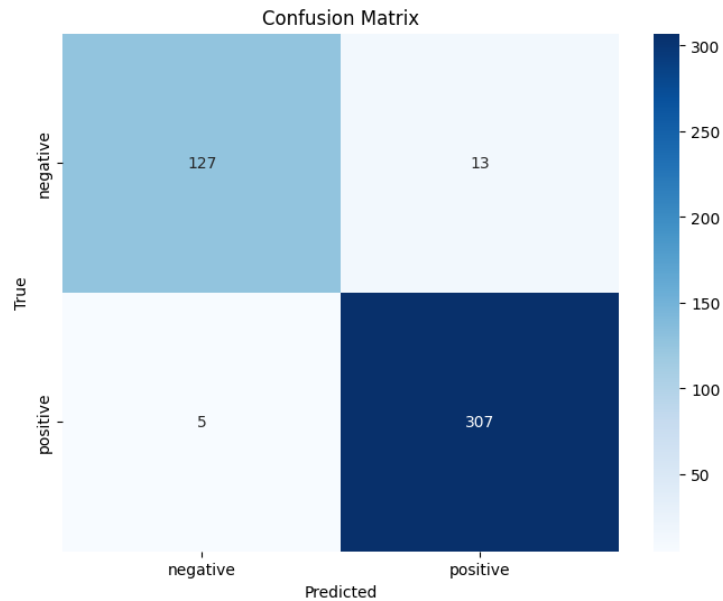
6. Hasil Evaluasi Akhir Menggunakan Data Test

Setelah model dilatih menggunakan 80% data latih, evaluasi dilakukan pada 20% data test yang telah dipisahkan sebelumnya untuk mengukur kemampuan model dalam mengklasifikasikan sentimen pada data yang tidak terlihat selama pelatihan. Hasil evaluasi pada data test memberikan metrik sebagai berikut:

- Akurasi: 96.61%
- Precision: 96.60%
- Recall: 96.61%
- F1-Score: 96.60%

Dengan Akurasi yang mencapai 96.61%, model menunjukkan kemampuan yang sangat baik dalam mengklasifikasikan tweet menjadi kategori positif atau negatif. Precision dan Recall yang tinggi (kedua-duanya di atas 96%) menunjukkan bahwa model mampu dengan baik mengklasifikasikan data positif dan negatif tanpa banyak kesalahan. F1-Score, yang menggabungkan precision dan recall, juga menunjukkan keseimbangan yang baik dalam klasifikasi.

Berdasarkan evaluasi kinerja model klasifikasi menggunakan confusion matrix, hasil yang diperoleh menunjukkan performa model yang sangat baik. Dari 332 data uji, model mampu mengklasifikasikan 307 data positif dengan benar sebagai positif (True Positive) dan 127 data negatif dengan benar sebagai negatif (True Negative). Tingginya nilai True Positive dan True Negative ini menunjukkan bahwa model memiliki tingkat akurasi yang tinggi dalam mengklasifikasikan data sesuai dengan kelas aslinya.



Gambar 5 confusion matrix

Matrix ini menunjukkan bahwa:

- **True Positives (TP):** 307 data positif diklasifikasikan dengan benar sebagai positif.
- **True Negatives (TN):** 127 data negatif diklasifikasikan dengan benar sebagai negatif.
- **False Positives (FP):** 13 data negatif yang salah diklasifikasikan sebagai positif.
- **False Negatives (FN):** 5 data positif yang salah diklasifikasikan sebagai negatif.

Dengan nilai True Positives yang sangat tinggi, model ini sangat baik dalam mengklasifikasikan sentimen positif dan negatif. Namun, ada beberapa False Positives dan False Negatives, meskipun jumlahnya sangat kecil. [1]

Berdasarkan hasil yang diperoleh, model GCN (Graph Convolutional Network) berhasil mencapai Akurasi sebesar 96.61%, yang menunjukkan kemampuan yang sangat baik dalam mengklasifikasikan sentimen pada data test. Dengan Precision dan Recall yang sangat tinggi, model ini juga tidak hanya mengklasifikasikan sebagian besar data dengan benar, tetapi juga mampu mengurangi jumlah False Negatives dan False Positives.

Confusion Matrix mengungkapkan bahwa model ini bekerja sangat baik dalam memisahkan kelas positif dan negatif, meskipun ada beberapa kasus di mana data negatif salah diklasifikasikan sebagai positif dan sebaliknya. Hasil ini menunjukkan bahwa model GCN dapat menangkap relasi antar kata dalam graf dengan sangat baik, mengklasifikasikan sentimen berdasarkan pola yang dipelajari selama pelatihan.

Meskipun hasil ini sudah sangat baik, masih ada potensi untuk lebih meningkatkan kinerja model, terutama dalam menangani False Positives dan False Negatives. Penelitian selanjutnya dapat mengeksplorasi penerapan model graf lain, seperti Graph Attention Networks (GAT), yang lebih fokus pada perhatian model terhadap node yang lebih relevan dalam graf.

D. Kesimpulan

Penelitian ini menunjukkan bahwa model Graph Convolutional Network (GCN) berhasil mengklasifikasikan sentimen publik terhadap isu penangkapan ikan ilegal oleh kapal Vietnam dengan akurasi mencapai 96.61%. Hasil evaluasi menggunakan precision, recall, dan F1-Score juga menunjukkan kinerja yang sangat baik, dengan precision sebesar 96.60% dan recall sebesar 96.61%, yang mencerminkan kemampuan model dalam membedakan dengan akurat antara kelas positif dan negatif. Precision yang tinggi menunjukkan bahwa sebagian besar prediksi positif yang dibuat oleh model adalah benar, sementara recall yang tinggi menandakan bahwa model berhasil mendeteksi sebagian besar data positif yang ada dalam dataset. Selain itu, F1-Score yang mencapai 96.60% menunjukkan keseimbangan yang baik antara precision dan recall, menggambarkan bahwa model ini tidak hanya akurat, tetapi juga efektif dalam mendeteksi dan mengklasifikasikan kedua kelas sentimen secara seimbang.

Hasil ini menegaskan bahwa model GCN sangat efektif dalam analisis sentimen berbasis graf, dengan kemampuannya untuk membedakan antara sentimen positif dan negatif secara akurat. Dalam konteks penelitian ini, kata-kata yang berhubungan dengan sentimen positif, seperti "tindak" dan "bertindak", sering muncul dalam diskusi mengenai kebijakan dan tindakan positif yang diambil terkait isu penangkapan ikan ilegal. Sementara itu, kata-kata yang berkaitan dengan sentimen negatif, seperti "illegal" dan "merugikan", dominan dalam diskusi mengenai pelanggaran hukum dan dampak buruk dari kegiatan ilegal tersebut. Visualisasi graf, seperti Word Cloud, memberikan wawasan tambahan mengenai kata-kata dominan yang berkontribusi pada analisis sentimen yang lebih mendalam, dengan kata-kata penting yang mewakili kedua kutub sentimen ini.

Secara keseluruhan, penelitian ini membuktikan bahwa GCN dapat digunakan sebagai metode yang sangat efektif untuk analisis sentimen berbasis graf. Model ini tidak hanya mampu memberikan hasil yang akurat dalam mengklasifikasikan sentimen, tetapi juga mampu memberikan gambaran yang lebih jelas tentang keterkaitan kata dalam sentimen positif dan negatif. Parameter yang digunakan dalam model ini meliputi input channels sebesar 1, output channels 2 untuk klasifikasi sentimen positif dan negatif, serta lapisan konvolusi pertama dengan 16 unit tersembunyi, diikuti oleh lapisan konvolusi kedua yang menghasilkan dua kelas sentimen. Optimizer yang digunakan adalah Adam dengan learning rate 0.01, dan fungsi kerugian yang digunakan adalah CrossEntropyLoss. Hasil penelitian ini membuka peluang untuk penelitian lebih lanjut, yang dapat mengadopsi data dari berbagai platform media sosial seperti Facebook, Instagram, atau platform lainnya, untuk meningkatkan keberagaman dan representasi data yang lebih komprehensif. Selain itu, pendekatan multibahasa yang melibatkan berbagai bahasa dapat digunakan untuk memperluas cakupan analisis sentimen dan mencakup opini yang lebih beragam dari masyarakat di berbagai negara.

Dengan hasil yang diperoleh, penelitian ini menunjukkan potensi besar Graph Convolutional Networks (GCN) dalam analisis sentimen berbasis graf, yang dapat diterapkan pada berbagai masalah sosial dan ekonomi, serta memberikan manfaat yang signifikan dalam pengambilan keputusan yang lebih informasional dan berbasis data.

E. Referensi

- [1] K. I. Pratiwi, N. R. Putri, and A. Efridadewi, "Aufklarung : Jurnal Pendidikan , Sosial dan Humaniora Analisis Penangkapan Ikan Ilegal oleh Kapal Asing (Vietnam) di Laut Natuna," vol. 3, no. 4, pp. 9–16, 2023.
- [2] A. B. Adhywidya and A. I. Budianto, "Upaya Hukum Terhadap Illegal Fishing Kapal Penangkap Ikan Vietnam Di ZEEI Legal Remedy against Illegal Fishing Vietnamese Fishing Vessels in Zeei," vol. 5, no. 2, pp. 247–256, 2023.
- [3] "No Title," 2022.
- [4] S. Nasional, T. Elektro, S. Informasi, and T. Informatika, "Analisis Percakapan di Media Sosial Twitter Terkait Pemindahan Ibu Kota Menggunakan Social Network Analysis Berbasis Model Jejaring Tersentralisasi," pp. 401–408, 2022.
- [5] A. L. Efendi, A. Fadilla, A. C. Khoirunnisa, G. N. Bakry, and N. Aristi, "Analisis Jaringan Komunikasi # Pilpres2024 Pada Platform Twitter," no. 204, pp. 219–232, 2024, doi: 10.32509/wacana.v22i2.2976.
- [6] N. Permatasari, R. Yosral, and C. F. Annisa, "(Twitter Analysis About Online Education During COVID-19 Pandemic In Indonesia) Analisis Media Sosial Twitter Tentang Pendidikan Daring Pada Masa Pandemi Covid- 19 Di Indonesia," pp. 359–369, 2020.
- [7] P. Studi, T. Industri, F. T. Industri, and U. I. Indonesia, "Analisis Strategi Marketing Online Travel Agent Pada Konten Media Sosial Twitter Menggunakan Sentiment Analysis Dan Social Network Analysis (Studi Kasus : Traveloka) ," 2023.
- [8] L. Anggraini and P. Matematika, "Penerapan metode matematika dalam analisis data big data," vol. 3, no. 1, pp. 1–13, 2023.
- [9] L. Pertiwi, "Penerapan Algoritma Text Mining , Steaming Dan Texrank Dalam Peringkasan Bahasa Inggris," vol. 1, no. 3, pp. 100–104, 2022.
- [10] N. Andrew *et al.*, "Implementasi Metode Word2Vec Dan TextRank Dalam Aplikasi Mobile Peringkasan Berita Olahraga," 2023.
- [11] I. G. A. Wibawa, "Analisis Sentimen Berbasis Aspek Ulasan Pelanggan Hotel di Bali Analisis Sentimen Berbasis Aspek Ulasan Pelanggan Hotel di Bali Menggunakan Metode Decision Tree," no. July 2022, 2023, doi: 10.24843/JLK.2023.v11.i03.p19.
- [12] M. Purnama, "Deep Learning Untuk Pembelajaran Representasi Graf : Metode Dan Aplikasi," vol. 4, no. 1, pp. 1–19, 2024.
- [13] S. Mazya, P. Tyas, R. Sarno, and B. S. Rintyarna, "Analisis Perbandingan Metode Klasifikasi Sentimen Berita Saham : Pendekatan Machine Learning , Deep Learning , Transfer Learning , dan Graf," vol. 9, no. 1, pp. 58–64, 2024.
- [14] N. Faulina *et al.*, "Enhancing Tuberculosis Diagnosis:Effective Naive Bayes Classification using SMOTE and Tomek Links for Imbalanced Data," vol. 6, no. 2, pp. 1–14, 2024.