

KOMPARASI METODE LABEL POWERSSET K-NN DAN ML-KNN DALAM KLASIFIKASI MULTI-LABEL CYBERBULLYING PADA KOMENTAR INSTAGRAM

Imamah Nur Fadlilah, Eka Dyar Wahyuni, Reisa Permatasari
Sistem Informasi, Universitas Pembangunan Nasional "Veteran" Jawa Timur
Jl. Raya Rungkut Madya, Gunung Anyar, Surabaya, Indonesia
imamahnf123@gmail.com

ABSTRAK

Peningkatan kasus *cyberbullying* di media sosial, khususnya pada aplikasi Instagram dengan berbagai bentuknya dapat berdampak negatif pada korban. Fenomena ini menunjukkan perlunya metode yang efektif untuk mengidentifikasi dan mengatasi komentar berisi *cyberbullying*. Penelitian ini bertujuan untuk mendeteksi komentar yang mengandung *cyberbullying* menggunakan pendekatan klasifikasi multi-label. Metode yang digunakan memanfaatkan ekstraksi fitur teks menggunakan TF-IDF dengan kombinasi *unigram*, *bigram* dan *trigram*, serta dua pendekatan utama dalam klasifikasi multi-label yaitu *Problem Transformation* dengan Label Powerset KNN dan *Algorithm Adaptation* dengan ML-KNN. Data yang digunakan dalam penelitian ini diambil dari komentar sebuah postingan pada aplikasi Instagram, dan akan dibangun sebanyak enam model klasifikasi *cyberbullying*. Hasil Penelitian menunjukkan bahwa pendekatan ML-KNN dengan kombinasi *unigram* memiliki akurasi yang lebih tinggi dibandingkan Label Powerset KNN dalam mendeteksi komentar yang mengandung *cyberbullying*.

Kata kunci : klasifikasi, multi-label, *cyberbullying*, LP-KNN, ML-KNN

1. PENDAHULUAN

Pada tahun 2024, Instagram menempati peringkat pertama sebagai platform media sosial paling banyak digunakan di Indonesia, dengan sekitar 84,8% pengguna internet aktif [1]. Platform ini digunakan untuk berbagi foto, video, dan cerita, serta berinteraksi melalui fitur komentar. Namun fitur komentar sering disalahgunakan untuk menyebarkan komentar negatif seperti penghinaan dan ujaran kebencian. Menurut Broadband Search, 2024, Instagram menempati peringkat pertama sebagai platform dengan insiden *cyberbullying* tertinggi, dengan persentase 42%. Hal ini menunjukkan bahwa meskipun fitur komentar Instagram memungkinkan interaksi yang positif, pengguna juga berisiko terkena *cyberbullying*.

Cyberbullying adalah tindakan perundungan yang dilakukan melalui teknologi digital termasuk media sosial, aplikasi chatting dan forum online, dengan tujuan menakut-nakuti, membuat marah, atau mempermalukan individu lain [2]. Terdapat beberapa bentuk *cyberbullying* yang sering muncul di media sosial, yaitu Flaming, Denigration, Harassment, dan Body Shaming. Cyberbullying memiliki dampak yang sangat merugikan, terutama bagi korban yang sering kali mengalami tekanan psikologis dan sosial. Di Indonesia, jumlah kasus *cyberbullying* telah meningkat secara signifikan dalam beberapa tahun terakhir. Data dari Komisi Perlindungan Anak Indonesia (KPAI) menunjukkan bahwa dari tahun 2011-2019, terdapat 2.473 laporan *cyberbullying*, dan jumlah ini terus meningkat setiap tahunnya [3].

Berbagai jenis *cyberbullying* dalam satu komentar sering muncul dalam kombinasi, sehingga klasifikasi multi-label memungkinkan setiap komentar dikategorikan dengan baik. Klasifikasi multi-label

adalah pendekatan yang digunakan untuk menangani kasus di mana satu sampel data dapat memiliki dua atau lebih label [4]. Ada dua pendekatan utama dalam klasifikasi multi-label, yaitu Problem Transformation dan Algorithm Adaptation. Problem Transformation, seperti Label Powerset (LP), mengubah masalah multi-label menjadi beberapa masalah klasifikasi *single-label*. Sebaliknya, Algorithm Adaptation mengubah algoritma *single-label* agar dapat langsung menangani masalah multi-label, seperti pada ML-KNN [5].

Berdasarkan latar belakang yang telah dijelaskan, penelitian ini berfokus pada penerapan klasifikasi multi-label dalam mendeteksi *cyberbullying* di Instagram. Penelitian ini membandingkan dua pendekatan utama dalam klasifikasi multi-label, yaitu Problem Transformation menggunakan Label Powerset KNN dan Algorithm Adaptation, yakni ML-KNN. Penelitian ini menggunakan teknik ekstraksi fitur TF-IDF untuk mengubah teks menjadi data terstruktur, sehingga algoritma dapat memprosesnya dengan akurat. Dengan membandingkan kedua pendekatan ini, penelitian ini bertujuan untuk menemukan metode yang paling akurat dalam mendeteksi berbagai bentuk *cyberbullying* di Instagram.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Beberapa penelitian sebelumnya telah menerapkan klasifikasi multi-label dalam berbagai bidang. Penelitian oleh [6] mengkategorikan fenomena ekonomi menggunakan pendekatan multi-label dengan algoritma seperti Naïve Bayes dan Decision Trees, di mana Label Powerset dengan Naïve Bayes menunjukkan performa terbaik dengan akurasi

63,22%. Sementara itu, penelitian oleh [7] menggunakan Multi-Label K-Nearest Neighbors (ML-KNN) untuk mengklasifikasikan pengguna media sosial berdasarkan topik yang mereka ikuti, menghasilkan akurasi 91% yang lebih unggul dibandingkan Binary Relevance dan Label Powerset. Penelitian lain [8] menerapkan klasifikasi multi-label pada Hadis Bukhari terjemahan Bahasa Indonesia dengan algoritma KNN yang dikombinasikan dengan Mutual Information untuk seleksi fitur, mencapai akurasi 91,14% dengan hamming loss 0,0886. Hasil penelitian ini menunjukkan bahwa berbagai metode klasifikasi multi-label memiliki keunggulan masing-masing tergantung pada karakteristik data yang digunakan.

2.2. TF-IDF

Metode TF-IDF merupakan suatu cara untuk memberikan bobot hubungan suatu kata (term) terhadap dokumen. Metode ini menggabungkan dua konsep untuk perhitungan bobot, yaitu frekuensi kemunculan sebuah kata di dalam sebuah dokumen tertentu dan inverse frekuensi dokumen yang mengandung kata tersebut. Frekuensi kemunculan kata di dalam dokumen yang diberikan menunjukkan seberapa penting kata itu di dalam dokumen tersebut. Sehingga bobot hubungan antara sebuah kata dan sebuah dokumen akan tinggi apabila frekuensi kata tersebut tinggi di dalam dokumen dan frekuensi keseluruhan dokumen yang mengandung kata tersebut yang rendah pada kumpulan dokumen [9].

2.3. K-NN

K-Nearest Neighbors (KNN) adalah salah satu metode yang paling sering digunakan dalam klasifikasi teks. KNN bekerja dengan mengklasifikasikan suatu objek berdasarkan sampel data latih terdekat dalam ruang fitur, guna menentukan kelompok dari k objek terdekat [10]. Metode ini menentukan klasifikasi dengan mengukur jarak terdekat antara objek baru dan kelompok data yang sudah ada. Jarak tersebut dihitung berdasarkan kemiripan antara data input dan data dalam kelompok, menggunakan sejumlah fitur yang relevan.

2.4. Label Powerset

Secara umum, metode klasifikasi multi-label dibagi menjadi dua kategori, yaitu metode transformasi data multi-label dan metode adaptasi algoritma multi-label. Salah satu pendekatan yang paling sederhana dalam transformasi data adalah Label Powerset (LP), yang menciptakan kelas baru untuk setiap kombinasi label dan kemudian menyelesaikan masalah ini menggunakan pendekatan klasifikasi multi-class [4]. Label Powerset mengubah langsung Kumpulan data multi-label menjadi Kumpulan data single-label dengan mentransformasikan setiap kombinasi label yang berbeda dalam data multi-label menjadi kelas yang berbeda dalam data single-label [11].

2.5. Multi-label KNN

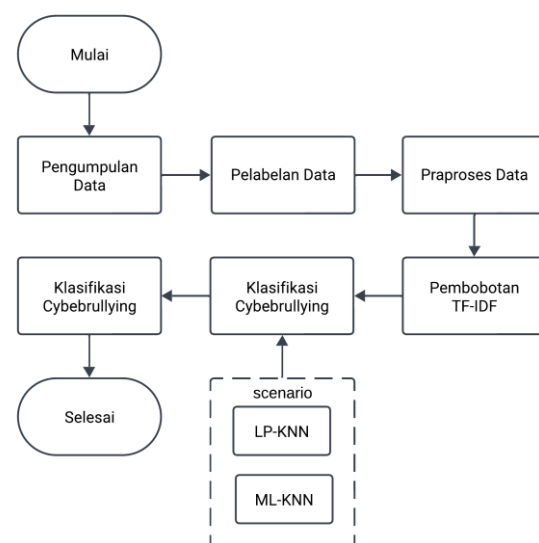
Multilabel K-Nearest Neighbors merupakan adaptasi dari algoritma K-Nearest Neighbors yang dirancang untuk menangani kasus dengan label ganda. Algoritma ini dikembangkan oleh [12] dan menjadi metode "lazy learning" pertama untuk masalah multilabel. Langkah awal dalam algoritma ini melibatkan perhitungan prior probabilities dan posterior probabilities, yang dilakukan hanya berdasarkan data latih. Implementasi KNN dilakukan pada tahap distribusi probabilitas posterior, dan prediksi probabilitas dilakukan dengan mengikuti prinsip aturan Bayesian.

2.6. Confusion Matrix

Confusion matrix atau matriks klasifikasi merupakan alat penting dalam menganalisis performa model klasifikasi. Matriks ini digunakan untuk mengevaluasi sejauh mana model mampu melakukan klasifikasi dengan akurat [13]. Confusion matrix membandingkan hasil prediksi model dengan nilai sebenarnya dari data uji. Dengan *confusion matrix*, kita dapat menghitung berbagai matriks evaluasi seperti *accuracy*, *precision*, *recall*, dan *F1-score*.

3. METODE PENELITIAN

Alur penelitian yang dilakukan ditunjukkan pada gambar diagram alir penelitian di bawah.



Gambar 1. Diagram alir penelitian

Pada setiap proses dalam Gambar 1 akan dijelaskan dalam poin-poin sebagai berikut

3.1. Pengumpulan Data

Data yang digunakan dalam penelitian ini adalah komentar-komentar yang diambil dari pengguna Instagram, yang diperoleh melalui teknik scraping dari beberapa akun dalam rentang tahun 2024. Data komentar kemudian disimpan dalam format CSV,

yang memudahkan proses analisis dan pengolahan data lebih lanjut.

3.2. Pelabelan Data

Pelabelan data dalam penelitian ini menggunakan lima label, yaitu *flaming*, *denigration*, *harassment*, *body shaming*, dan *neutral*. Proses dilakukan secara multi-label, di mana satu komentar bisa memiliki lebih dari satu label. Semua proses pelabelan dilakukan secara manual dan dalam setiap tahap pelabelan, label dominan dari setiap komentar akan diambil untuk digunakan pada tahap analisis berikutnya.

3.3. Praproses Data

Data yang telah diberi label harus melalui tahapan praproses sebelum ke tahapan selanjutnya. Tujuan tahapan ini untuk membuat data lebih terstruktur dan konsisten sehingga model yang akan dibangun dapat bekerja dengan lebih optimal..

3.4. Pembobotan dengan TF-IDF

Tahap selanjutnya, data latih akan diproses menggunakan TF-IDF untuk memberi bobot pada setiap kata berdasarkan kepentingannya dalam teks [14]. Selain itu, akan digunakan n-gram (1-3) sebagai skenario percobaan, di mana model akan mempertimbangkan kombinasi kata dalam satuan unigram, bigram, dan trigram untuk menangkap lebih banyak informasi dalam teks..

3.5. Klasifikasi Cyberbullying

Klasifikasi dilakukan dengan menggunakan metode LP-KNN dan ML-KNN. Model ini dilatih menggunakan data latih yang telah diberi label dan dibobot sebelumnya. Dalam skenario ini, teks akan dikategorikan ke dalam lima label utama, yaitu *Neutral* (tidak mengandung *cyberbullying*), *Flaming*, *Denigration*, *Harassment*, dan *Body Shaming*..

3.6. Evaluasi Model

Tahap ini bertujuan untuk menentukan model terbaik berdasarkan berbagai skenario yang telah diuji. Evaluasi dilakukan dengan membandingkan performa model LP-KNN dan ML-KNN pada setiap kombinasi skenario. Evaluasi ini mencakup pengukuran *accuracy*, *precision*, *recall*, *F1-score*, serta *hamming loss*.

4. HASIL DAN PEMBAHASAN

4.1. Pengumpulan Data

Data dalam penelitian ini dikumpulkan dari komentar di akun *public figure* yang menjadi sorotan di Instagram sepanjang tahun 2024. Akun-akun ini dipilih berdasarkan berbagai isu yang menarik perhatian publik dan memicu banyak perbincangan di kalangan netizen. Proses pengumpulan data dilakukan dengan *scraping* menggunakan Jupyter Notebook dan *library* Selenium untuk mengambil komentar secara otomatis.

	username	text
0	amahsnh	Lah jd nya cinlok ? Difilm sakinah kaka ini ce...
1	syairazull	Salsa di katin selingkuhan si iki lo klarifik...
2	beby.memey	Pemeran di sinetronnya jadi real 🤔
3	silvmint	Kasian jadi ga bisa spill bebeb 🤔
4	missamoey69	Oh ini yg jdi selingkuhan
5	who_aamm_iii1	Jangan ya dek ya, jangan ambil milik orang lain 🤔
6	fulloflovemoment	udah kayak ibu ibu anak 2 mukanya alias tuaaa...
7	verdiaspp_	gelayy.
8	chalila_f	Mau bersin tapi gajadi ya ka???
9	pangsit.nyonyor	Kalo muka kurenggg minimal hati baik

Gambar 2. Hasil pengumpulan data

Hasil scraping pada Gambar 2 menghasilkan total 6.455 data komentar yang tersimpan dalam dua kolom, yaitu kolom *username* yang menunjukkan akun pembuat komentar dan kolom *text* yang berisi isi komentar. Data yang diperoleh kemudian diproses lebih lanjut untuk keperluan analisis klasifikasi.

4.2. Pelabelan Data

Proses pelabelan data dilakukan secara manual oleh lima annotator, di mana setiap komentar dapat diberikan lebih dari satu label sesuai dengan kategori yang telah ditentukan, yaitu *Flaming*, *Denigration*, *Harassment*, *Body Shaming*, dan *Neutral*. Karena penelitian ini menggunakan klasifikasi multi-label, satu komentar dapat termasuk dalam lebih dari satu kategori jika mengandung unsur yang relevan.

	username	text	label1_pipit	label1_aka	label1_4	label1_5	label1_6	label1_7
0	israanaya	nyta lu pndh den beben gue	[Denigration]	[Denigration]	[Denigration]	[Denigration]	[Denigration]	[Denigration]
1	_palmabada11	Makasih ya dlmna, kamu udh ngertiin Riky, Sika	[Denigration]	[Flaming]	[Body Shaming]	[Denigration]	[Denigration]	[Denigration]
2	alya.x	Kasian apik klo punya cewk kandel ngendog ga b...	[Flaming]	[Denigration]	[Denigration]	[Denigration]	[Denigration]	[Denigration]
3	amptugan	Kasih juch sama cewa 🤔	[Denigration]	[Flaming]	[Body Shaming]	[Denigration]	[Flaming]	[Denigration]
4	malakyy	Makasi Uda nntah Rik kanyaa cewa sama si kndh	[Denigration]	[Denigration]	[Denigration]	[Denigration]	[Denigration]	[Denigration]

Gambar 3. Hasil pelabelan data

Setiap annotator pada Gambar 3 secara independen memberikan label pada komentar yang telah dikumpulkan dari Instagram. Setelah seluruh komentar dilabeli, label yang paling dominan berdasarkan kesepakatan dari kelima annotator akan ditetapkan sebagai label akhir untuk memastikan konsistensi dalam klasifikasi data.

4.3. Praproses Data

Tahapan praproses data dilakukan untuk meningkatkan kualitas teks sebelum masuk ke tahap pemodelan. Proses ini mencakup beberapa langkah utama, yaitu *cleaning*, *case folding*, normalisasi, penghapusan *stopwords*, *stemming*, dan tokenisasi. *Cleaning* dilakukan untuk menghapus elemen yang tidak relevan, seperti mention, hashtag, URL, emoji, angka, dan simbol lainnya. *Case folding* kemudian mengubah semua huruf menjadi kecil agar konsistensi data terjaga. Selanjutnya, normalisasi mengubah kata tidak baku atau slang ke bentuk formal menggunakan kamus yang telah disusun sebelumnya.

Setelah normalisasi, penghapusan *stopwords* dilakukan untuk menghilangkan kata-kata umum yang tidak memiliki makna signifikan dalam analisis. Stemming kemudian digunakan untuk mengubah kata ke bentuk dasarnya agar variasi kata tetap dianggap sama. Terakhir, tokenisasi memecah teks menjadi daftar kata-kata terpisah yang siap untuk diproses lebih lanjut. Hasil dari seluruh tahap praproses ini disimpan selanjutnya pada Tabel 1 akan digunakan sebagai input utama dalam perancangan model klasifikasi.

Tabel 1. Contoh hasil dari tahap praproses

Text	Praproses Text
Ngerasa dosa banget abis makan salad	['rasa', 'dosa', 'sangat', 'habis', 'makan', 'salad']
G seimbang sma kpla	['tidak', 'imbang', 'sama', 'kepala']
Nggak sudi banget dengar lagunya sama kenal orgnya	['tidak', 'sudi', 'sangat', 'dengar', 'lagu', 'sama', 'kenal', 'orang']

4.4. Pembobotan dengan TF-IDF

Tahap selanjutnya adalah pembobotan kata menggunakan TF-IDF yang dikombinasikan dengan n-grams (1-3) untuk menangkap pola kata dalam teks. Tahapan ini bertujuan agar kata-kata dalam data dapat direpresentasikan dalam bentuk numerik yang dapat dikenali oleh model.

Hasil dari metode pembobotan kata TF-IDF yang dikombinasikan dengan n-grams (1-3) pada tahap ini, akan diterapkan pada setiap skenario pemodelan menggunakan algoritma Label Powerset K-Nearest Neighbors (LP-KNN) dan Multi-label K-Nearest Neighbors (ML-KNN).

Tabel 2. Jenis skenario pemodelan

Nama Model	Deskripsi
Model 1	LP-KNN dengan TF-IDF Unigram
Model 2	LP-KNN dengan TF-IDF Bigram
Model 3	LP-KNN dengan TF-IDF Trigram
Model 4	ML-KNN dengan TF-IDF Unigram
Model 5	ML-KNN dengan TF-IDF Bigram
Model 6	ML-KNN dengan TF-IDF Trigram

Tabel 2 menunjukkan skenario pemodelan yang melibatkan kombinasi algoritma LP-KNN dan ML-KNN dengan metode TF-IDF menggunakan Unigram, Bigram, dan Trigram. Pemodelan diuji skenario dengan pembagian data 70:30.

4.5. Klasifikasi Cyberbullying

Tahapan ini model mengklasifikasikan komentar ke dalam kategori *Flaming*, *Denigration*, *Harassment*, *Body Shaming*, atau *Neutral*. Pendekatan ini bertujuan untuk mengevaluasi sejauh mana model mampu mengenali serta membedakan berbagai jenis *cyberbullying* berdasarkan pola dan karakteristik teks yang telah diproses sebelumnya.

Tabel 3. Hasil klasifikasi cyberbullying

Nama Model	Accuracy	Precision	Recall	F1-Score	Hamming Loss
Model 1	0.56	0.81	0.55	0.56	0.20
Model 2	0.48	0.79	0.45	0.43	0.24
Model 3	0.18	0.68	0.42	0.30	0.42
Model 4	0.64	0.71	0.77	0.74	0.13
Model 5	0.60	0.80	0.63	0.64	0.17
Model 6	0.12	0.51	0.44	0.40	0.24

Tabel 3 merupakan hasil klasifikasi menggunakan dua pendekatan utama, yaitu LP-KNN dan ML-KNN, dengan variasi teknik ekstraksi fitur TF-IDF Unigram, Bigram, dan Trigram. Hasil ini menunjukkan perbandingan performa masing-masing model dalam mengenali dan membedakan kategori *cyberbullying*. Selanjutnya, model-model tersebut akan dievaluasi lebih lanjut untuk menentukan metode terbaik yang dapat digunakan dalam mendeteksi berbagai bentuk *cyberbullying*.

4.6. Evaluasi Model

Berdasarkan hasil evaluasi performa model, ML-KNN dengan TF-IDF Unigram (Model 4) menunjukkan hasil terbaik dibandingkan model lainnya. Model ini memiliki *accuracy* tertinggi (0.64), *F1-score* tertinggi (0.74), serta *Hamming Loss* terendah (0.13). Dibandingkan dengan pendekatan LP-KNN, model ML-KNN secara umum memiliki performa yang lebih baik, terutama dalam menangkap pola klasifikasi multilabel. Penggunaan TF-IDF Unigram juga terbukti lebih optimal dibandingkan Bigram dan Trigram, yang cenderung mengalami penurunan performa pada kedua pendekatan algoritma.

LP-KNN dengan TF-IDF Unigram (Model 1) menjadi model terbaik dalam kelompok LP-KNN dengan *accuracy* 0.56 dan *F1-score* 0.56, tetapi masih kalah dibandingkan model ML-KNN. Sementara itu, penggunaan Bigram dan Trigram dalam kedua pendekatan justru menyebabkan penurunan performa yang cukup signifikan. Model LP-KNN dengan TF-IDF Trigram (Model 3) dan ML-KNN dengan TF-IDF Trigram (Model 6) memiliki *accuracy* terendah, masing-masing sebesar 0.18 dan 0.12, menunjukkan bahwa kombinasi ini kurang efektif dalam klasifikasi *multi-label* pada dataset yang digunakan.

Dengan mempertimbangkan seluruh metrik evaluasi, ML-KNN dengan TF-IDF Unigram (Model 4) merupakan model terbaik dalam penelitian ini. Model ini memiliki keseimbangan yang baik antara *precision*, *recall*, dan *F1-score*, serta tingkat kesalahan yang lebih rendah dibandingkan model lainnya.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil evaluasi, dapat disimpulkan bahwa ML-KNN menunjukkan performa yang lebih baik dibandingkan Label Powerset KNN dalam klasifikasi *multi-label* untuk mendeteksi *cyberbullying* di Instagram. Model terbaik dalam penelitian ini adalah ML-KNN dengan TF-IDF Unigram, yang

memiliki *accuracy* tertinggi (0.64), F1-score tertinggi (0.74), serta *Hamming Loss* terendah (0.13). Hasil ini menunjukkan bahwa pendekatan ML-KNN lebih efektif dalam menangani klasifikasi *multi-label* dibandingkan LP-KNN, serta penggunaan Unigram dalam teknik ekstraksi fitur memberikan representasi teks yang lebih optimal dibandingkan Bigram dan Trigram.

Penelitian berikutnya dapat mempertimbangkan penggunaan metode *word embedding* seperti Word2Vec atau FastText untuk meningkatkan pemahaman model terhadap konteks bahasa dalam komentar Instagram. Dengan hasil yang diperoleh, penelitian ini memberikan kontribusi dalam memahami berbagai pendekatan klasifikasi *multi-label* dalam mendeteksi *cyberbullying*. Implementasi model yang lebih optimal diharapkan dapat digunakan untuk membangun sistem deteksi otomatis yang lebih akurat dalam mengidentifikasi berbagai bentuk *cyberbullying* di platform media sosial, khususnya Instagram.

DAFTAR PUSTAKA

- [1] R. H. Pandjaitan, "The Social Media Marketing Mix Trends in Indonesia for 2024: Communication Perspective," *Jurnal Komunikasi Ikatan Sarjana Komunikasi Indonesia*, vol. 9, no. 1, pp. 251–269, Jun. 2024, doi: 10.25008/jkiski.v9i1.1005.
- [2] G. W. Giumetti and R. M. Kowalski, "Cyberbullying via social media and well-being," *Curr Opin Psychol*, vol. 45, p. 101314, Jun. 2022, doi: 10.1016/j.copsyc.2022.101314.
- [3] N. Dewi, "Pentingnya Perlindungan terhadap Korban Cyberbullying di Indonesia." [Online]. Available: https://tirto.id/pentingnya-perlindungan-terhadap-korban-cyberbullying-di-indonesia-f5io#google_vignette
- [4] J. Bogatinovski, L. Todorovski, S. Džeroski, and D. Kocev, "Comprehensive comparative study of multi-label classification methods," *Expert Syst Appl*, vol. 203, p. 117215, Oct. 2022, doi: 10.1016/j.eswa.2022.117215.
- [5] N. Endut, W. M. A. F. W. Hamzah, I. Ismail, M. Kamir Yusof, Y. Abu Baker, and H. Yusoff, "A Systematic Literature Review on Multi-Label Classification based on Machine Learning Algorithms," *TEM Journal*, pp. 658–666, May 2022, doi: 10.18421/TEM112-20.
- [6] Nofriani and N. Budi Kurniawan, "Harnessing Multi-label Classification Approaches for Economic Phenomena Categorization," *ASEAN Journal on Science and Technology for Development*, vol. 38, no. 2, Aug. 2021, doi: 10.29037/ajstd.680.
- [7] A. Huang, R. Xu, Y. Chen, and M. Guo, "Research on multi-label user classification of social media based on ML-KNN algorithm," *Technol Forecast Soc Change*, vol. 188, p. 122271, Mar. 2023, doi: 10.1016/j.techfore.2022.122271.
- [8] A. Hanafi, A. Adiwijaya, and W. Astuti, "Klasifikasi Multi Label pada Hadis Bukhari Terjemahan Bahasa Indonesia Menggunakan Mutual Information dan k-Nearest Neighbor," *Jurnal Sisfokom (Sistem Informasi dan Komputer)*, vol. 9, no. 3, pp. 357–364, Sep. 2020, doi: 10.32736/sisfokom.v9i3.980.
- [9] H. J. Meijer, J. Truong, and R. Karimi, "Document Embedding for Scientific Articles: Efficacy of Word Embeddings vs TFIDF," 2021, [Online]. Available: <http://arxiv.org/abs/2107.05151>
- [10] J. Sreemathy and P. Balamurugan, "AN EFFICIENT TEXT CLASSIFICATION USING KNN AND NAIVE BAYESIAN," 2012. [Online]. Available: <https://api.semanticscholar.org/CorpusID:15986279>
- [11] S. Kumar, N. Kumar, A. Dev, and S. Naorem, "Movie genre classification using binary relevance, label powerset, and machine learning classifiers," *Multimed Tools Appl*, vol. 82, no. 1, pp. 945–968, Jan. 2023, doi: 10.1007/s11042-022-13211-5.
- [12] G. Tian, J. Wang, R. Wang, G. Zhao, and C. He, "A multi-label social short text classification method based on contrastive learning and improved <scp>ml-KNN</scp>," *Expert Syst*, vol. 41, no. 7, Jul. 2024, doi: 10.1111/exsy.13547.
- [13] M. Heydarian, T. E. Doyle, and R. Samavi, "MLCM: Multi-Label Confusion Matrix," *IEEE Access*, vol. 10, pp. 19083–19095, 2022, doi: 10.1109/ACCESS.2022.3151048.
- [14] D. Zeng, E. Zha, J. Kuang, and Y. Shen, "Multi-label text classification based on semantic-sensitive graph convolutional network," *Knowl Based Syst*, vol. 284, p. 111303, Jan. 2024, doi: 10.1016/j.knosys.2023.111303