

Analisis Sentimen Universitas Kristen Immanuel Menggunakan SVM, GAN dan SMOTE

Jatmika¹, Liefson Jacobus², Devant Joe Raffael³, Edys Kuswanto⁴

Informatika, Universitas Kristen Immanuel, Yogyakarta

Jln Ukrim No 23 Cupuwatu Purwomartani, Sleman, Yogyakarta

Email : jatmika@ukrimuniversity.ac.id

Abstrak

Penelitian ini bertujuan untuk melakukan analisis sentimen terhadap Universitas Kristen Immanuel (UKRIM) dengan menggunakan metode klasifikasi Support Vector Machine (SVM). Data yang dikumpulkan dari media sosial melalui proses web scrapping selanjutnya diproses melalui beberapa tahapan preprocessing seperti pembersihan data, tokenisasi, dan penghapusan stopword. Salah satu tantangan utama dalam penelitian ini adalah ketidakseimbangan kelas dalam data, yang diatasi dengan menggunakan dua teknik, yaitu Synthetic Minority Over-sampling Technique (SMOTE) dan Generative Adversarial Network (GAN). Teknik SMOTE digunakan untuk melakukan oversampling terhadap kelas minoritas, sedangkan GAN digunakan untuk menghasilkan data artifisial yang menyerupai data asli. Untuk proses ekstraksi fitur, digunakan pendekatan berbasis Bidirectional Encoder Representations from Transformers (BERT), yang kemudian diklasifikasikan menggunakan model SVM. Hasil evaluasi dari eksperimen menunjukkan bahwa penggunaan GAN untuk membuat data buatan dan SMOTE untuk menyeimbangkan kelas-kelas pada data meningkatkan performa model klasifikasi, karena model SVM secara efektif dapat membedakan sentimen negatif, netral, dan positif.

Kata kunci: SVM, SMOTE, GAN, analisis sentimen, media sosial

Abstract

Social media has become an important tool for people to express their opinions about various institutions, including universities. This research aims to analyze Universitas Kristen Immanuel (UKRIM) using the Support Vector Machine (SVM) classification method. Data collected from social media through web scraping was processed through several preprocessing stages such as data cleaning, tokenization, and stopword removal. One of the main challenges in this study is class imbalance, which was addressed using two techniques: Synthetic Minority Over-sampling Technique (SMOTE) and Generative Adversarial Network (GAN). SMOTE was used to oversample the minority class, while GAN was utilized to generate artificial data that resembles real input. Feature extraction was carried out using the Bidirectional Encoder Representation from Transformers (BERT) model, and classification was performed using SVM. The evaluation results of the experiments show that the use of GAN to create artificial data and SMOTE to balance the classes in the data improves the classification model performance, as the SVM model can effectively distinguish negative, neutral, and positive sentiments.

Keywords: SVM, SMOTE, GAN, sentiment analysis, social media

1. PENDAHULUAN

Pada era digital, media sosial menjadi sumber informasi bagi masyarakat sekaligus wadah untuk memberikan opini serta ulasan mengenai berbagai layanan dan institusi. Untuk sebuah institusi pendidikan seperti Universitas Kristen Immanuel, memahami persepsi publik menjadi aspek yang penting dalam pengambilan keputusan strategis. Opini publik yang tersebar di platform media sosial dapat memberikan wawasan berharga dalam pengambilan keputusan yang bisa meningkatkan reputasi serta pelayanan dari institusi.

Platform seperti Google Maps dan X (sebelumnya Twitter) bisa menjadi sumber data yang sangat berpotensi dalam melakukan analisis sentimen. Platform Google Maps memungkinkan pengguna untuk memberikan ulasan secara langsung terhadap suatu institusi, sehingga opini masing-masing pengguna bisa diakses secara terbuka. Sementara itu, X juga bisa menjadi sumber data yang ideal dalam melakukan analisis sentimen karena kebebasan beropini yang ditawarkan platform ini terhadap penggunaannya, sehingga data yang



dikeluarkan oleh platform ini bisa sangat bermanfaat untuk melakukan monitoring opini publik serta analisis sentimen media sosial [1].

Namun, melakukan analisis sentimen menggunakan data dari media sosial akan menimbulkan beberapa tantangan, yaitu kurangnya jumlah data yang sesuai dengan konteks serta ketidakseimbangan sentimen yang muncul. Konten yang diambil dari media sosial biasanya memiliki distribusi yang tidak seimbang; pada satu topik yang sama, salah satu kelas sentimen mungkin lebih dominan dibanding kelas sentimen yang lain. Ketidakseimbangan ini bisa mempengaruhi tingkat akurasi model analisis sentimen, di mana model tersebut akan menjadi bias terhadap satu kelas sentimen dibanding yang lain [2]. Selain itu, walaupun media sosial memberikan data dalam jumlah yang sangat besar, mencari data yang relevan dengan suatu institusi tertentu bisa menjadi masalah tersendiri, karena tidak semua topik mendapat jumlah perhatian yang sama dari pengguna media sosial. Sehingga jumlah data yang diambil mungkin tidak cukup untuk memberikan wawasan yang berarti bagi model [3][4].

Menurut Birjali, analisis sentimen merupakan suatu tugas yang bertujuan untuk mengekstrak dan menganalisis pendapat, sentimen, sikap, dan persepsi masyarakat terhadap berbagai entitas, seperti topik, produk, dan layanan [5]. Analisis sentimen telah diterapkan di berbagai negara untuk mengevaluasi institusi pendidikan tinggi. Di Maroko, analisis sentimen digunakan untuk memprediksi sentimen di media sosial secara real-time, yang dapat membantu kementerian pendidikan dalam pengambilan keputusan [6]. Analisis sentimen membantu universitas dalam meningkatkan kualitas pengajaran dan fasilitas dengan mengidentifikasi area yang memerlukan perbaikan berdasarkan umpan balik mahasiswa [7][8].

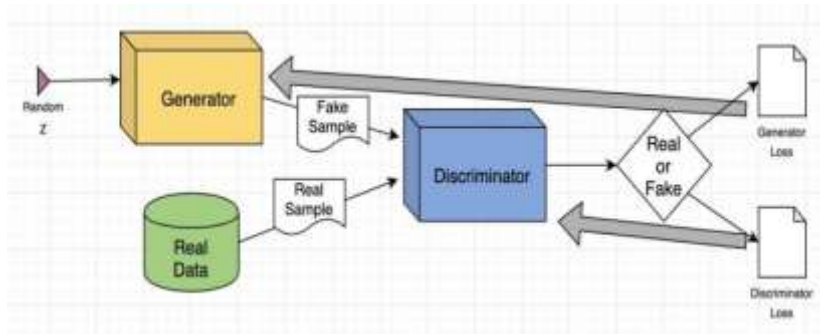
Perkembangan platform media sosial dan ulasan seperti Twitter dan Google Maps memungkinkan analisis opini publik yang menyeluruh. Dengan penggunaan media sosial yang ramai, Twitter sangat berguna untuk analisis sentimen, sementara Google Maps memberikan kemudahan dalam mencari informasi mengenai lokasi suatu objek serta ulasan dari pengunjung yang dapat dianalisis untuk memahami sentimen mereka [9]. Gambaran yang lebih lengkap tentang analisis sentimen dapat diperoleh dengan menggabungkan data dari kedua platform tersebut. Dalam penelitian ini, digunakan tiga metode utama untuk analisis sentimen, yaitu Support Vector Machine (SVM), Generative Adversarial Network (GAN), dan Synthetic Minority Oversampling Technique (SMOTE).

Dalam analisis sentiment, mesin vektor pendukung (SVM) adalah teknik pembelajaran mesin canggih yang dapat menangani klasifikasi dan regresi masalah besar dengan jumlah data pelatihan yang lebih sedikit [10]. SVM menunjukkan tingkat akurasi yang signifikan dan sangat efektif dalam mendeteksi polaritas teks dan menangani data yang sangat besar. Misalnya, penggunaan SVM dalam analisis sentimen pada ulasan produk dan media sosial menghasilkan tingkat akurasi yang tinggi sebesar 91,8% dalam analisis tweet tentang maskapai penerbangan AS [11].

Penelitian tentang analisis sentimen rokok yang membandingkan penggunaan metode SVM dan Naive Bayes menunjukkan bahwa metode SVM memiliki keakuratan yang lebih tinggi dalam mengklasifikasikan polaritas teks (positif atau negatif). Ini membuka pandangan baru tentang opini publik, yang mungkin bermanfaat bagi pembuat kebijakan [12].

Generative Adversarial Networks (GANs) adalah model generatif kecerdasan buatan yang digunakan untuk mengidentifikasi distribusi probabilitas data pelatihan dan kemudian menggunakan distribusi tersebut untuk menghasilkan data baru yang realistis. Generator dan discriminator adalah dua komponen utama GAN, yang dilatih secara bersamaan melalui metode pelatihan adversarial [13]. Pada kasus sentimen analisis twitter dengan menggunakan GAN yang dilakukan oleh Mahalakshmi, penggunaan GAN dalam analisis sentimen Twitter menunjukkan bahwa klasifikasi sentimen lebih akurat. Metode ini mencapai akurasi klasifikasi sebesar 93.33% dengan memanfaatkan Convolutional Neural Network (CNN) untuk ekstraksi fitur, mengungguli metode sebelumnya dalam hal akurasi, recall, ketepatan, dan skor F1 [14].





Gambar 1. Visualisasi cara kerja GAN

GAN terdiri dari generator yang membuat data baru dan discriminator yang mengevaluasi keaslian data tersebut. Keduanya terlibat dalam permainan minimax, di mana generator berusaha menipu discriminator dengan data palsu yang tampak nyata, sementara discriminator berusaha membedakan antara data nyata dan palsu. GAN dilatih secara berlawanan. Baik generator maupun discriminator berusaha untuk meningkatkan kemampuan mereka dalam menghasilkan data yang realistis [15].

Dalam analisis sentimen, ketidakseimbangan data adalah masalah umum. Ini terjadi ketika satu kelas (misalnya, sentimen positif) memiliki dominasi yang lebih besar daripada yang lain. Sehingga model tidak bias terhadap kelas mayoritas, SMOTE membantu mengatasi masalah ini dengan mengumpulkan data sintesis dari kelas minoritas [16].

SMOTE menghasilkan sampel sintesis untuk kelas minoritas dalam set data dengan membuat titik data baru secara artifisial. SMOTE membantu menyeimbangkan distribusi kelas, yang sangat penting untuk meningkatkan kinerja algoritma klasifikasi pada set data tidak seimbang [17].

Dengan menyeimbangkan distribusi kelas data, SMOTE terbukti meningkatkan akurasi model analisis sentimen. Misalnya, penggunaan SMOTE dengan SVM menghasilkan akurasi tertinggi sebesar 94,38% dalam analisis sentimen terhadap ulasan aplikasi pemerintah [18].

Berdasarkan penelitian yang telah dilakukan oleh Arief Negara, penggunaan SMOTE dalam analisis sentimen terbukti efektif dalam meningkatkan akurasi dan kemampuan model untuk mendeteksi kelas minoritas. Dengan menyeimbangkan data, SMOTE membantu mengatasi bias yang disebabkan oleh ketidakseimbangan kelas, sehingga menghasilkan model yang lebih andal dan akurat dalam berbagai konteks analisis sentimen [19].

2. METODOLOGI PENELITIAN

Pada bagian ini akan dibahas metodologi penelitian terkait pengumpulan data, pemrosesan yang dilakukan terhadap data, serta pelatihan model yang digunakan.

2.1. Pengambilan Data

Data yang akan digunakan dalam penelitian ini diambil menggunakan aplikasi third-party, yaitu ScrapingDog dan Apify. Aplikasi ScrapingDog digunakan untuk mengambil seluruh data ulasan Universitas Kristen Immanuel dari Google Maps, sedangkan Apify digunakan untuk mengambil 100 unggahan paling baru pada aplikasi X yang mengandung kata-kata “ukrim”, atau “universitas kristen immanuel”, atau “ukrim jogja”.

Cara kerja dua aplikasi ini dalam mengambil data cukup berbeda. Aplikasi Apify mengharuskan penggunanya untuk mengakses halaman website tersebut, memilih media sosial yang ingin diambil data unggahannya (dalam kasus ini adalah X), memasukkan keyword dan saringan-saringan lainnya yang diinginkan, lalu memulai proses scraping. Hasil dari proses ini adalah file .json yang berisi data yang telah diminta.

Sedangkan ScrapingDog memerlukan script bahasa Python untuk mengambil data. Pengguna akan menggunakan kunci API tersebut untuk membuat sebuah request terhadap endpoint milik ScrapingDog yang sesuai dengan platform yang ingin diambil datanya. Respons dari request ini adalah sebuah file .json yang bisa langsung diubah menjadi dictionary, dan bisa langsung dipakai atau di-unggah untuk dipakai di lain waktu.

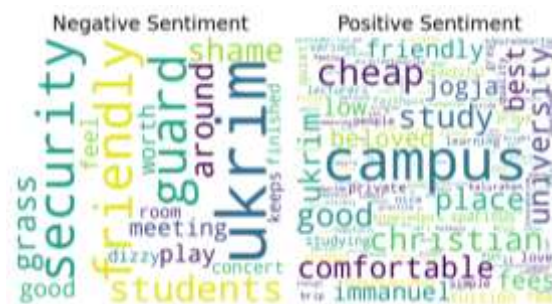
Setelah mengambil data, dilakukan labeling terhadap masing-masing data secara manual untuk menghitung akurasi dari masing-masing model yang akan dicoba.

2.2. Pra-pemrosesan Data

Data yang didapat dari masing-masing aplikasi akan dibersihkan terlebih dahulu, kolom-kolom seperti ID, tanggal, dan user akan dibuang dan dijadikan dalam bentuk objek pandas DataFrame. Setelah itu, kedua objek DataFrame ini akan digabung menjadi satu buah DataFrame yang terdiri dari dua kolom, yaitu:

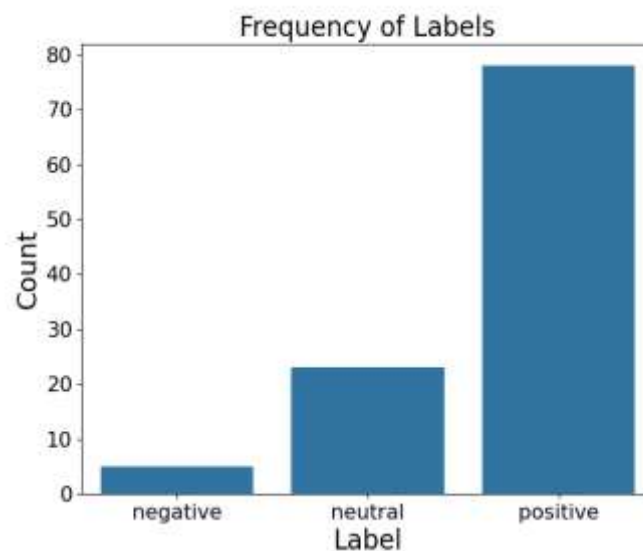
- 1) snippet: berisi teks yang akan diproses oleh model.
- 2) label: berisi label yang telah dimasukkan secara manual. Berisi salah satu dari 3 kelas “negative”, “neutral”, atau “positive”.

Setelah itu, kolom snippet akan diproses supaya merapikan baris-baris yang memiliki teks yang berulang-ulang, contohnya teks “Good CampusGood Campus” akan dirapikan menjadi “Good Campus” saja. Perulangan ini terjadi akibat penerjemahan otomatis yang dilakukan oleh Google. Cara merapikannya adalah dengan membagi teks menjadi dua bagian yang sama panjangnya, jika bagian satu sama dengan bagian dua, maka bagian kedua akan dihapus menyisakan bagian satu saja. Hasil dari proses ini bisa dilihat melalui grafik di bawah ini, di mana semakin sering sebuah kata muncul, maka ukuran kata tersebut di dalam grafik akan menjadi semakin besar.



Gambar 2. Visualisasi *WordCloud* data teks yang diambil

Seluruh pemrosesan di atas menghasilkan data dengan frekuensi masing-masing label sebagai berikut:



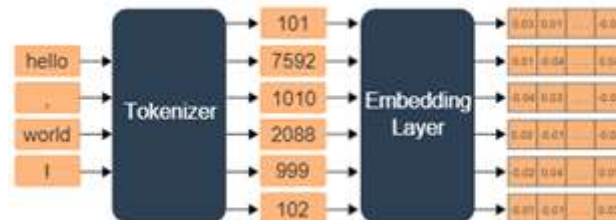
Gambar 3. Grafik persebaran frekuensi label teks data dari *Google Maps* dan *X*

Dari grafik tersebut bisa dilihat bahwa terjadi ketidakseimbangan yang sangat besar pada data yang telah diambil. Kelas positive memiliki jumlah lebih dari 70 data, neutral hanya sekitar 20 data, dan kelas negative bahkan tidak mencapai 10. Selain itu, jumlah data secara keseluruhan yang sangat sedikit dalam konteks analisis sentimen akan sangat mempengaruhi kemampuan model untuk mengambil ilmu dan memprediksi sentimen dari teks.

Setelah ini, data akan dipisah berdasarkan rasio 70:30, dimana 70% dari data akan dipakai untuk proses pelatihan model, dan sisa 30% dari data akan dipakai untuk proses pengujian akurasi model.

2.3. Proses Embedding

Data yang masih berbentuk teks harus diubah menjadi bentuk angka. Proses mengubah data teks menjadi angka disebut dengan Embedding. Proses embedding ini akan merubah satu baris teks menjadi sebuah vektor numerik yang merepresentasikan fitur-fitur dari teks tersebut (biasa disebut dengan vektorisasi). Pada penelitian ini, akan digunakan BERT-Embedding karena kemampuannya dalam melakukan vektorisasi terhadap emoji dan tanda baca secara langsung.



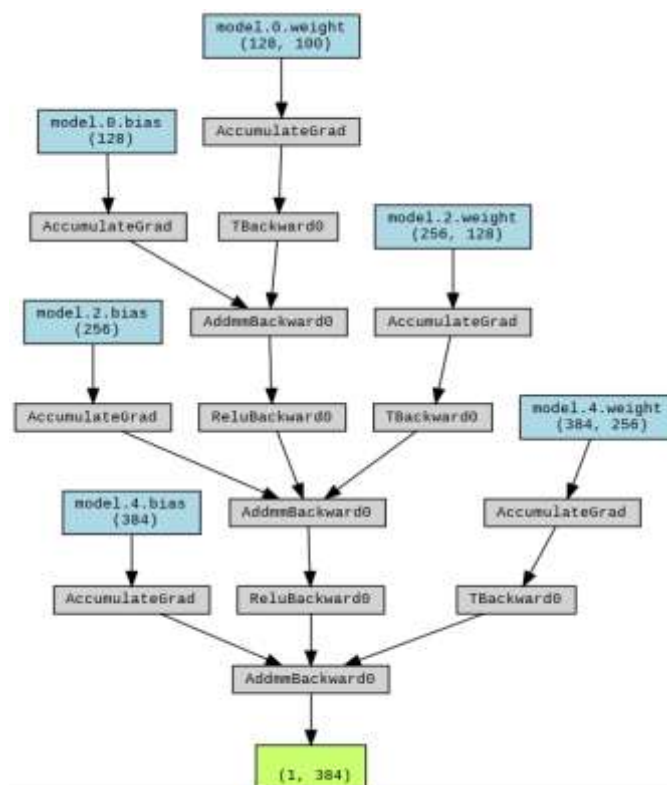
Gambar 4. Visualisasi cara kerja *BERT*

Hal ini perlu dilakukan karena sebagian besar model Machine Learning tidak bisa memproses data teks secara mentah, dan perlu diubah menjadi sebuah vektor angka yang bisa merepresentasikan teks tersebut secara keseluruhan, sehingga model bisa melakukan proses pelatihan dan prediksinya dengan maksimal.

2.4. Arsitektur GAN

Jaringan GAN memerlukan dua buah model yang akan bekerja berlawanan untuk menghasilkan data buatan yang realistis. Model Generator akan membuat data artifisial, sedangkan Discriminator akan membandingkan hasil data buatan yang telah dibuat dengan data asli, dan mengklasifikasikan apakah data artifisial tersebut sudah cukup mirip dengan data asli. Model Generator akan terus menerus membuat data buatan hingga Discriminator menganggap data yang dibuat sudah cukup mirip dengan data asli. Berikut adalah arsitektur dari masing-masing model tersebut:

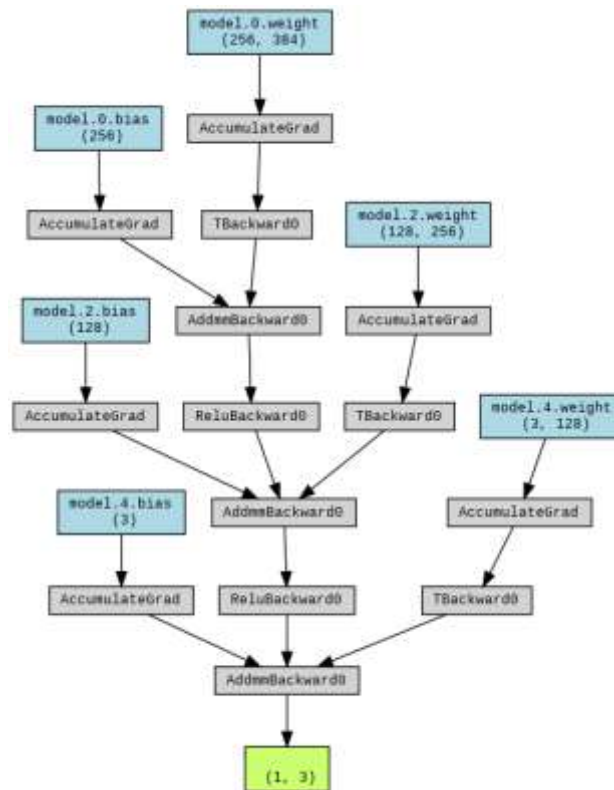
1) Model Generator



Gambar 5. Arsitektur model *generator*

Model Generator ini terdiri dari 5 layer, yaitu 3 layer linear biasa, dan 2 layer ReLu di antara masing-masing layer linear. Model ini dioptimasi menggunakan Adam, dan menggunakan CrossEntropy sebagai loss functionnya. Model ini menerima 100 angka random noise sebagai input, dan mengeluarkan vektor satu dimensi berukuran (1, 384) sebagai data artifisial yang siap dinilai oleh Discriminator, dan hasil penilaian tersebut akan digunakan untuk backtracking dan melatih model Generator.

2) Model Discriminator



Gambar 6. Arsitektur model *discriminator*

Model Discriminator ini terdiri dari 5 layer, yaitu 3 layer linear biasa, dan 2 layer ReLu di antara masing-masing layer linear. Model ini dioptimasi menggunakan Adam, dan menggunakan CrossEntropy sebagai loss functionnya. Model ini dilatih dengan cara memprediksi kelas data asli dan data buatan dari Generator, lalu menggunakan rata-rata loss dari kedua prediksi kelas tersebut untuk melakukan backtracking.

2.5. Penyeimbangan Kelas Menggunakan SMOTE

Penyeimbangan data menggunakan SMOTE dilakukan hingga data-data dengan kelas minoritas (seperti neutral, dan negative) memiliki jumlah data yang sama dengan kelas mayoritas. Jadi, jika kelas positive memiliki 55 jumlah data, maka SMOTE akan terus melakukan oversampling hingga kelas neutral dan negative memiliki jumlah data yang sama dengan kelas positive.

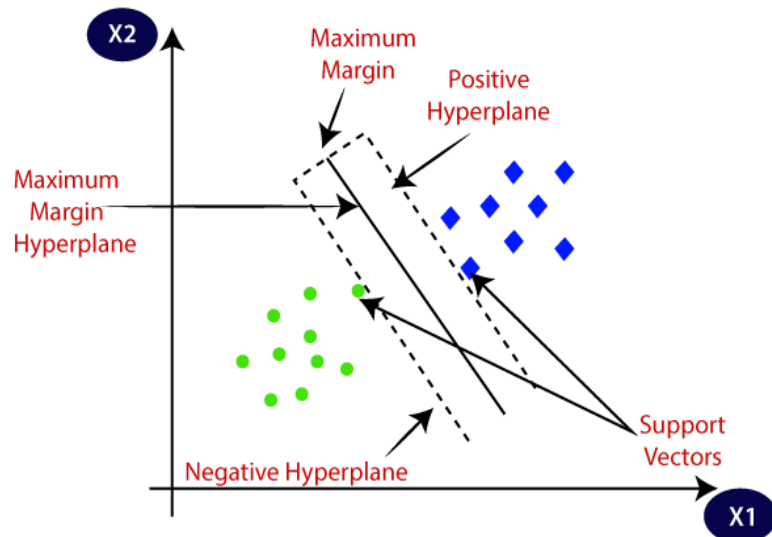
Cara SMOTE melakukan oversampling adalah dengan melakukan interpolasi antar data-data dengan kelas yang sama. Interpolasi ini dilakukan dengan mengambil salah satu data dalam kelas yang minoritas, mencari x titik data yang paling dekat dengan titik tersebut (nilai x adalah bilangan bulat yang bisa diubah-ubah), dan menciptakan sampel baru yang segaris dengan titik awal dan titik data yang terdekat tersebut [22][23].

2.6. Proses Pelatihan Model

Model yang akan digunakan pada penelitian ini adalah Support Vector Machine (SVM). Model ini digunakan karena karakteristiknya yang cocok digunakan saat suatu data memiliki dimensionalitas (jumlah fitur) yang tinggi. Hal ini sangat cocok karena setelah dilakukan embedding, data yang digunakan akan memiliki lebih dari 300 fitur.

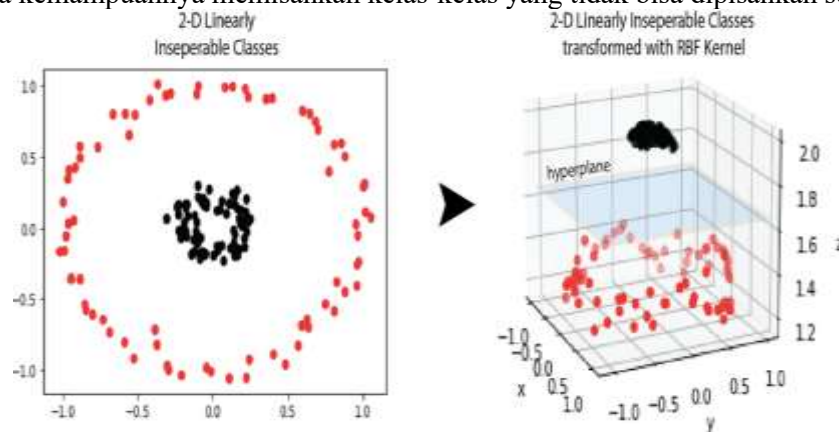
SVM menggunakan fungsi kernel untuk memetakan input data ke dalam dimensi data yang lebih tinggi, hal ini memungkinkan SVM untuk mencari titik-titik pemisahan data yang tidak bisa dipisahkan secara linear pada dimensi sebelumnya yang lebih rendah. Model SVM memiliki target utama yaitu mencari hyperplane

yang memaksimalkan jarak antara hyperplane tersebut dan masing-masing kelas yang ada di dalam data. Jarak inilah yang disebut dengan support vector. Model SVM hanya menggunakan support vector yang dibuat untuk menentukan hyperplane yang akan dibuat, sehingga membuat SVM cukup efisien dalam segi memori dan beban komputasi [20][21].



Gambar 7. Ilustrasi cara kerja SVM pada data 2 dimensi

Maka dari itu, penentuan kernel sangat penting dalam penggunaan model SVM, karena kernel yang berbeda bisa mempengaruhi bentuk dan lokasi hyperplane yang dibuat oleh model. Penelitian ini akan menggunakan kernel “rbf” karena kemampuannya memisahkan kelas-kelas yang tidak bisa dipisahkan secara linear biasa.



Gambar 8. Ilustrasi hasil kerja kernel RBF pada 3 dimensi

Dalam penelitian ini akan dilakukan tiga kali proses pelatihan model, yaitu:

- 1) Model SVM saja,
- 2) Model SVM menggunakan data hasil SMOTE untuk penyeimbangan kelas,
- 3) Model SVM menggunakan data hasil augmentasi GAN yang telah diseimbangkan menggunakan SMOTE.

Hasil dari ketiga percobaan inilah yang akan dibandingkan dan dievaluasi pada penelitian ini.

3. HASIL DAN PEMBAHASAN

3.

Berdasarkan pengambilan data melalui ScrapingDog dan Apify diperoleh total 106 data dari Google Maps dan X. Data tersebut kemudian diberikan label dan diperoleh pengelompokan seperti gambar berikut dengan label positif berjumlah 78 data, 23 data berlabel netral, dan 5 data berlabel negatif. Kinerja model dalam klasifikasi dapat dipengaruhi oleh ketidakseimbangan ini, terutama untuk kelas minoritas.

Untuk mengatasi masalah ini, tiga skenario pelatihan model SVM dilakukan. Untuk menilai pengaruh penyeimbangan data terhadap kinerja model dalam mengklasifikasikan, tiga metode ini dibandingkan. Pertama, model SVM dilatih dengan data asli tanpa penyeimbangan; kedua, metode Oversampling Minority Synthetic (SMOTE) digunakan untuk melatih model SVM dengan data yang telah diseimbangkan; dan terakhir, metode SMOTE digunakan untuk melatih model SVM dengan data hasil augmentasi Generative Adversarial Network (GAN) yang kemudian diseimbangkan kembali.

3.1. Model SVM

Model SVM dilatih tanpa menggunakan penyesuaian kelas, data pelatihan dan data uji dibagi dengan rasio 70:30.

Tabel 1. Hasil Evaluasi Model SVM

Variabel	<i>Precision</i>	<i>Recall</i>	<i>F1-score</i>	<i>Support</i>
Negatif	0.00	0.00	0.00	2
Netral	0.00	0.00	0.00	7
Positif	0.72	1.00	0.84	23
Akurasi			0.72	32

Hasil pengujian menunjukkan bahwa model menghasilkan akurasi sebesar 72%, dan 9 dari 32 data uji, atau 28% dari total, salah diklasifikasikan. Semua data berlabel "netral" dan "negatif" tidak diklasifikasikan dengan benar; sebaliknya, semuanya diprediksi sebagai "positif", seperti yang ditunjukkan oleh nilai precision, recall, dan f1-score untuk label "netral" dan "negatif" yang masing-masing bernilai 0.

Sebaliknya, model menunjukkan kinerja yang baik pada label "positif" dengan f1-score sebesar 0,84, yang menunjukkan bahwa model sangat bias terhadap kelas mayoritas, yaitu "positif", yang memiliki jumlah yang jauh lebih besar daripada kedua kelas lainnya. Faktor utama yang menyebabkan kinerja model yang buruk terhadap kelas minoritas adalah ketidakseimbangan distribusi label.

3.2. Model SVM menggunakan data hasil SMOTE

Dilakukan proses oversampling menggunakan SMOTE, sehingga distribusi kelas dalam data latih menjadi seimbang, yaitu masing-masing kelas positif, netral, dan negatif memiliki 55 data. Hal ini memberikan model SVM kesempatan belajar yang adil dari seluruh kelas yang sebelumnya tidak seimbang, di mana kelas positif mendominasi.

Tabel 2. Hasil Evaluasi Model SVM dengan SMOTE

Variabel	<i>Precision</i>	<i>Recall</i>	<i>F1-score</i>	<i>Support</i>
Negatif	0.00	0.00	0.00	2
Netral	0.67	0.29	0.40	7
Positif	0.76	0.96	0.85	23
Akurasi			0.75	32

Hasil evaluasi menunjukkan bahwa model bekerja lebih baik daripada versi sebelumnya. Akurasi model meningkat menjadi 75% dengan 8 kesalahan prediksi dari 32 data (25%), sedikit lebih rendah dari versi sebelumnya yang mencapai 28.12%. Model juga menunjukkan skor f1-tertinggi pada kelas positif (0.85), tetapi masih cukup rendah pada kelas netral (0.40), dan tidak mengenali kelas negatif sama sekali (0.00). Hal ini menunjukkan bahwa meskipun SMOTE membantu memperbaiki keseimbangan data dan meningkatkan performa keseluruhan, model masih cenderung bias terhadap kelas mayoritas (positif), serta kesulitan dalam membedakan antara ekspresi netral dan positif. Beberapa teks netral seperti "my campus" atau "balance of faith and science" masih salah diklasifikasikan sebagai positif, menunjukkan bahwa representasi semantik dari model belum sepenuhnya sensitif terhadap nuansa kalimat.



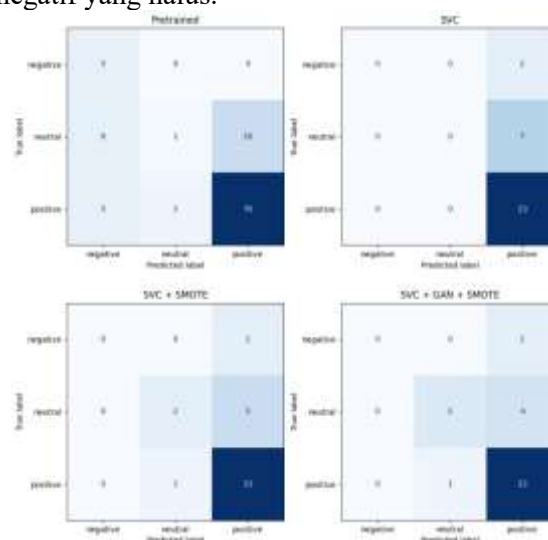
3.3. Model SVM menggunakan data hasil augmentasi GAN yang telah diseimbangkan menggunakan SMOTE.

Tabel 3. Hasil Evaluasi Model SVM dengan SMOTE dan GAN

Variabel	Precision	Recall	F1-score	Support
Negatif	0.00	0.00	0.00	2
Netral	0.75	0.43	0.55	7
Positif	0.79	0.96	0.86	23
Akurasi			0.78	32

Meskipun distribusi data telah diseimbangkan, model yang dilatih pada data hasil augmentasi menunjukkan akurasi sebesar 78% pada data uji. Kelas positif memiliki performa terbaik dengan F1-score sebesar 0.86, diikuti oleh kelas netral dengan F1-score sebesar 0.55, dan kelas negatif tidak terdeteksi sama sekali dengan F1-score sebesar 0.00. Ini menunjukkan bahwa, meskipun distribusi data telah diseimbangkan, model masih cenderung bias terhadap prediksi kelas positif.

Contoh misklasifikasi seperti "*my campus*" dan "*spirit of spirit*" menunjukkan kecenderungan model untuk memberikan label positif atau negatif pada teks netral atau bahkan yang memiliki nada ambigu atau makna yang tidak jelas. Selain itu, kritik terhadap fasilitas kampus juga dapat salah dianggap sebagai positif. Ini termasuk kalimat tentang rumput tinggi atau penjaga keamanan. Ini menunjukkan bahwa model kurang peka terhadap komentar netral dan negatif yang halus.



Gambar 9. Confusion matrix

4. SIMPULAN

Penelitian ini menunjukkan bahwa augmentasi data menggunakan GAN, serta penyeimbangan kelas data menggunakan SMOTE terbukti efektif dalam meningkatkan kinerja model dalam melakukan analisis sentimen terhadap Universitas Kristen Immanuel, terlihat dari nilai akurasinya yang tertinggi dibanding metode lain di nilai 78%, serta F1-score di kelas netral yang paling tinggi di 55%.

Namun semua model masih belum bisa mengenali kelas negatif dengan baik, hal ini mungkin terjadi karena jumlah sampel nyata kelas negatif yang sangat sedikit, dan model SVM yang belum bisa menyimpan dan memproses sentimen kalimat secara keseluruhan, sehingga kalimat yang mengandung kata seperti "not friendly" tidak dianggap negatif karena memiliki kata "friendly" di dalamnya.

Untuk penelitian selanjutnya, mungkin bisa menggunakan model yang bisa menyimpan informasi kalimat secara keseluruhan seperti BiLSTM atau BiGRU, serta melakukan hyperparameter tuning terhadap model-model yang dipakai. Pengaplikasian weighted loss function terhadap kelas-kelas yang tidak seimbang juga

bisa dipertimbangkan supaya model yang digunakan tidak mengalami bias terhadap salah satu kelas yang memiliki jumlah yang dominan.

DAFTAR PUSTAKA

- [1]. B. Sangam & S. Sangam, "An Ensemble model to analyze the sentiments of Textual Data for Monitoring and Comprehending Racist Views", *International Conference on Advances in Science and Technology (ICAST)*, vol. 6, hal. 187-190, 2023.
- [2]. Q. A. Xu, V. Chang, dan C. Jayne, "A systematic review of social media-based sentiment analysis: Emerging trends and challenges", *Decision Analytics Journal*, vol. 3, hal. 100073, 2022.
- [3]. K. Chakraborty, S. Bhattacharyya and R. Bag, "A Survey of Sentiment Analysis from Social Media Data," *IEEE Transactions on Computational Social Systems*, vol. 7, no. 2, hal. 450-464, April 2020.
- [4]. M. Rodríguez-Ibáñez, A. Casáñez-Ventura, F. Castejón-Mateos, P-M. Cuenca-Jiménez, "A review on sentiment analysis from social media platforms", *Expert Systems with Applications*, vol. 223, hal.119862, Maret 2023
- [5]. M. Birjali, M. Kasri, dan A. Beni-Hssane, "A comprehensive survey on sentiment analysis: Approaches, challenges and trends", *Knowledge-Based Systems*, vol. 226, hal. 107134, 2021.
- [6]. I. Lasri, A. Riadsolh, and M. Elbelkacemi, "Real-time Twitter Sentiment Analysis for Moroccan Universities using Machine Learning and Big Data Technologies," *International Journal of Emerging Technologies in Learning (iJET)*, vol. 18, no. 05, hal. 42–61, Maret. 2023
- [7]. N. T. P. Giang, T. T. Dien, and T. T. M. Khoa, "Sentiment analysis for university students' feedback," in *Advances in intelligent systems and computing*, 2020..
- [8]. O. F. Chamorro-Atalaya et al., "Identification of the satisfaction of university students through sentiment analysis: a systematic review," *International Journal of Evaluation and Research in Education (IJERE)*, vol. 14, no. 1, hal. 37, November 2024.
- [9]. A. D. Cahyani, "Analisa Kinerja Metode Support Vector Machine untuk Analisa Sentimen Ulasan Pengguna Google Maps," *Journal of Computer System and Informatics (JoSYC)*, vol. 4, no. 3, hal. 604–613, Mei 2023.
- [10]. A. Roy and S. Chakraborty, "Support vector machine in structural reliability analysis: A review," *Reliability Engineering & System Safety*, vol. 233, hal. 109126, Januari 2023.
- [11]. Q. Huang, "Sentiment analysis for social media using SVM classifier of machine learning," *Applied and Computational Engineering*, vol. 4, no. 1, pp. 86–90, Mei 2023.
- [12]. N. Tundo, R. Eldina, K. Setiawan, and R. Fajri, "Sentiment Analysis of Cigarette Use Based on Opinions from X Using Naive Bayes and SVM," *Jurnal Indonesia Manajemen Informatika Dan Komunikasi*, vol. 5, no. 3, hal. 2561–2569, September 2024.
- [13]. D. Saxena and J. Cao, "Generative Adversarial Networks (GANs)," *ACM Computing Surveys*, vol. 54, no. 3, hal. 1–42, Mei 2021.
- [14]. V. Mahalakshmi, P. Shenbagavalli, S. Raguvaran, V. Rajakumareswaran, and E. Sivaraman, "Twitter sentiment analysis using conditional generative adversarial network," *International Journal of Cognitive Computing in Engineering*, vol. 5, hal. 161–169, Januari. 2024.
- [15]. S. Islam et al., "Generative Adversarial Networks (GANs) in Medical Imaging: advancements, applications, and challenges," *IEEE Access*, vol. 12, hal. 35728–35753, Januari 2024.
- [16]. H. Yi, Q. Jiang, X. Yan, and B. Wang, "Imbalanced classification based on minority clustering synthetic minority oversampling technique with wind turbine fault detection application," *IEEE Transactions on Industrial Informatics*, vol. 17, no. 9, pp. 5867–5875, December. 2020.
- [17]. N. Sholihah, F. F. Abdulloh, M. Rahardi, A. Aminuddin, B. P. Asaddulloh, and A. Y. A. Nugraha, "Feature Selection Optimization for Sentiment Analysis of Tax Policy Using SMOTE



- and PSO,” International Conference on Smart Cities, Automation & Intelligent Computing Systems (ICON-SONICS), hal. 44–48, December 2023.
- [18]. M. Hadwan, M. Al-Sarem, F. Saeed, and M. A. Al-Hagery, “An Improved Sentiment Classification Approach for Measuring User Satisfaction toward Governmental Services’ Mobile Apps Using Machine Learning Methods with Feature Engineering and SMOTE Technique,” *Applied Sciences*, vol. 12, no. 11, hal. 5547, Mei 2022.
- [19]. A. B. P. Negara, “The influence of applying stopword removal and SMOTe on Indonesian sentiment classification,” *Lontar Komputer Jurnal Ilmiah Teknologi Informasi*, vol. 14, no. 3, hal. 172, Desember 2023.
- [20]. M. Shalaby, M. Farouk, & H. Khater, “Data reduction for SVM training using density-based border identification”, *PLOS ONE*, vol. 19, hal. 4, April 2024.
- [21]. S. Zhou, "Sparse SVM for Sufficient Data Reduction," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 44, no. 9, hal. 5560-5571, September 2022.
- [22]. S. Wang, Y. Dai, J. Shen, et al., “Research on expansion and classification of imbalanced data based on SMOTE algorithm,” *Sci Rep*, vol. 11, hal. 24039, Desember 2021.
- [23]. N. A. Azhar, M. S. M. Pozi, A. M. Din & A. Jatowt, "An Investigation of SMOTE Based Methods for Imbalanced Datasets With Data Complexity Analysis," *IEEE Transactions on Knowledge and Data Engineering*, vol. 35, no. 7, hal. 6651-6672, Juli 2023

