

ANALISIS SENTIMEN TERHADAP PROGRAM MSIB DALAM X 2024 MENGUNAKAN LEXICON-BASED DAN NAÏVE BAYES CLASSIFIER

Yuri Saputri, Betha Nurina Sari, Sofi Defiyanti

Teknik Informatika, Universitas Singaperbangsa Karawang

Jl. HS.Ronggo Waluyo, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat 41361, Telp. (0267) 641177
2110631170041@student.unsika.ac.id

ABSTRAK

Media sosial telah menjadi saluran utama bagi masyarakat dalam menyampaikan opini secara cepat, termasuk mengenai program pendidikan seperti Magang dan Studi Independen Bersertifikat (MSIB). Salah satu isu yang banyak dikeluhkan oleh peserta MSIB adalah keterlambatan pencairan Bantuan Biaya Hidup (BBH). Penelitian ini bertujuan untuk menganalisis sentimen mahasiswa terhadap program MSIB tahun 2024 melalui platform X (Twitter) dengan menggunakan metode Lexicon-Based dan Naïve Bayes Classifier. Proses analisis dilakukan melalui tahapan Knowledge Discovery in Database (KDD) yang mencakup data selection, preprocessing, transformation, data mining, dan evaluation. Data dikumpulkan menggunakan crawling tools dengan kombinasi kata kunci terkait MSIB dan BBH. Metode Lexicon-Based digunakan untuk pelabelan sentimen berdasarkan kamus sentimen, sedangkan Naïve Bayes digunakan untuk klasifikasi berdasarkan pembobotan TF-IDF. Hasil analisis menunjukkan bahwa mayoritas sentimen terhadap program MSIB bersifat negatif, terutama terkait isu keterlambatan pencairan dana. Model evaluasi menggunakan confusion matrix menunjukkan nilai akurasi terbaik diperoleh pada rasio pembagian data 70:30 dengan akurasi sebesar 87%, precision 90%, recall 88%, dan f-measure 86%. Temuan ini menunjukkan bahwa kombinasi metode Lexicon-Based dan Naïve Bayes mampu mengklasifikasikan sentimen dengan baik dan dapat digunakan sebagai masukan untuk perbaikan program MSIB ke depannya.

Kata kunci : Analisis Sentiment, Media Sosial, MSIB, *Lexicon-Based*, *Naïve Bayes Classifier*.

1. PENDAHULUAN

Perkembangan teknologi internet telah membawa dampak besar dalam berbagai aspek kehidupan, termasuk dunia pendidikan. Transformasi digital memungkinkan pembelajaran dilakukan secara fleksibel, tidak terbatas oleh ruang dan waktu. Mahasiswa kini dapat mengakses berbagai sumber daya pembelajaran melalui internet dan perangkat digital, menjadikan proses belajar lebih adaptif dan relevan dengan kebutuhan dunia kerja. Salah satu inovasi dalam pendidikan tinggi di Indonesia yang lahir dari transformasi ini adalah Program Magang dan Studi Independen Bersertifikat (MSIB). Program ini merupakan bagian dari kebijakan Kampus Merdeka yang bertujuan memberikan mahasiswa pengalaman langsung di dunia industri serta pelatihan bersertifikat. Program ini tidak hanya menjembatani dunia akademik dan dunia kerja, tetapi juga menyiapkan lulusan yang lebih kompeten dan siap bersaing di era digital [1].

Namun, meskipun MSIB menawarkan banyak manfaat, program ini juga menghadapi sejumlah permasalahan dalam pelaksanaannya, terutama terkait keterlambatan pencairan Bantuan Biaya Hidup (BBH) kepada mahasiswa peserta. Berbagai penelitian sebelumnya, seperti yang dilakukan [2] dan [3], menunjukkan bahwa sebagian besar mahasiswa mengalami keterlambatan pencairan dana, yang menyebabkan keresahan di kalangan peserta. Bahkan, banyak keluhan disampaikan melalui media sosial seperti Twitter (X), dan petisi daring telah dibuat sebagai bentuk protes terhadap ketidakjelasan

pencairan dana. Permasalahan ini tidak hanya mengganggu kenyamanan mahasiswa dalam mengikuti program, tetapi juga memunculkan pertanyaan mengenai transparansi dan efektivitas pengelolaan program tersebut. Oleh karena itu, perlu dilakukan kajian berbasis data untuk memahami lebih dalam persepsi mahasiswa terhadap program ini [4][5].

Penelitian ini bertujuan untuk melakukan analisis sentimen terhadap Program MSIB berdasarkan data yang diperoleh dari media sosial X (Twitter) selama tahun 2024. Secara khusus, penelitian ini memiliki dua tujuan utama, yaitu: pertama, menganalisis sentimen mahasiswa terhadap program MSIB, terutama dalam isu keterlambatan pencairan BBH, dengan mengelompokkan sentimen ke dalam kategori positif dan negatif menggunakan metode Lexicon-Based dan Naïve Bayes Classifier; kedua, mengukur tingkat akurasi dari kedua metode tersebut dalam mengklasifikasikan sentimen. Dengan demikian, penelitian ini diharapkan dapat memberikan gambaran yang objektif dan akurat tentang opini publik mahasiswa terhadap program MSIB.

Untuk mencapai tujuan tersebut, penelitian akan dilakukan melalui beberapa tahapan. Data dikumpulkan dari media sosial X dengan menggunakan kombinasi kata kunci dan hashtag yang relevan dengan Program MSIB. Setelah data terkumpul, dilakukan proses prapemrosesan seperti pembersihan teks, tokenisasi, dan normalisasi. Kemudian, data dianalisis menggunakan metode Lexicon-Based untuk pelabelan awal sentimen,

dilanjutkan dengan algoritma Naïve Bayes untuk klasifikasi yang lebih mendalam [6]. Model yang dihasilkan akan dievaluasi menggunakan metrik seperti akurasi, precision, recall, dan F1-score. Dengan pendekatan ini, diharapkan hasil penelitian dapat memberikan kontribusi nyata dalam memahami opini mahasiswa dan menjadi masukan bagi pihak terkait dalam memperbaiki program MSIB ke depan.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Dalam bahasa Indonesia, analisis sentimen adalah cara untuk menemukan bagaimana sebuah perasaan diungkapkan dalam teks dan diklasifikasikan menjadi sentimen positif, negatif, atau netral. Analisis sentimen bekerja secara sistematis untuk mengenali, mengekstrak, serta meneliti keadaan serta data subjektif. Analisis sentimen secara luas diterapkan pada analisis komentar konsumen, pembahasan, asumsi survei, serta media sosial. Tugas utama analisis sentimen merupakan untuk mengklasifikasikan polaritas data dalam dokumen baik kalimat ataupun kata-kata [7].

2.2. Naïve Bayes Classifier

Naïve Bayes adalah salah satu algoritma yang digunakan dalam teknik klasifikasi. Algoritma ini didasarkan pada prinsip yang dikembangkan oleh seorang ilmuwan asal Inggris bernama Thomas Bayes [8].

$$P(H|X) = \frac{P(X|H) \cdot P(H)}{P(X)} \quad (1)$$

Keterangan :

- X : Data sampel yang labelnya belum diketahui.
- H : Dugaan atau asumsi (hipotesis) bahwa data memiliki label tertentu.
- $P(H)$: Kemungkinan atau peluang bahwa hipotesis H itu benar.
- $P(X)$: Kemungkinan munculnya data sampel X secara umum.
- $P(X|H)$: Peluang munculnya data X jika memang hipotesis H itu benar.

$P(H|X)$: Peluang bahwa hipotesis H benar setelah melihat data sampel X

2.3. Lexicon Based

Lexicon-Based adalah salah satu cara umum dalam menganalisis sentimen pada media sosial. Metode ini menggunakan kamus sebagai sumber bahasa yang berfungsi untuk mengklasifikasikan sentimen dari setiap opini sehingga kalimat sentimen dapat diklasifikasikan ke dalam kelas negatif dan positif [9].

2.4. Penelitian sebelumnya

Penelitian sebelumnya yang dilakukan oleh An Nisaa et al., (2022) dalam jurnalnya yang berjudul "Performance of the Independent Campus Policy in Certified Internship and Independent Study Programs" menunjukkan bahwa keterlambatan pencairan uang

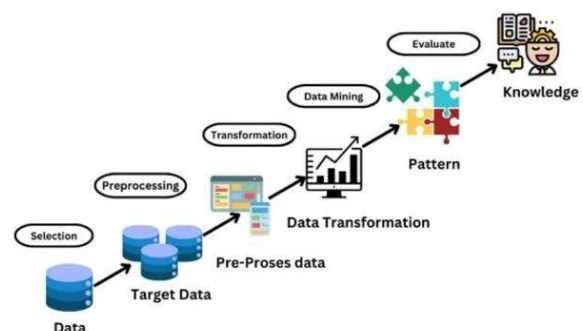
saku menjadi salah satu masalah yang sering dikeluhkan oleh mahasiswa peserta Magang dan Studi Independen Bersertifikat (MSIB). Berdasarkan hasil dari penelitian tersebut, banyak mahasiswa mengalami keterlambatan pencairan uang saku, bahkan ada yang belum menerima sama sekali dalam beberapa bulan. Masalah ini menimbulkan keresahan yang cukup besar, dibuktikan dengan adanya petisi yang dibuat oleh mahasiswa peserta MSIB batch pertama dan telah ditandatangani oleh 11.021 mahasiswa. Selain itu, muncul juga berbagai keluhan di media sosial dalam bentuk utas yang berisi tuntutan terkait pencairan dana yang belum terealisasi.

Penelitian lainnya yang dilakukan oleh Rahman et al., (2023) yang berjudul "Implementasi Kebijakan Pada Program Magang dan Studi Independen Bersertifikat Indonesia" menunjukkan bahwa keterlambatan pencairan BBH menjadi kendala utama dalam pelaksanaan program tersebut. Dalam jurnal tersebut terdapat hasil survei terhadap 8.402 mahasiswa peserta MSIB, ditemukan bahwa 92,05% mahasiswa mengalami keterlambatan pencairan BBH, sementara 74% mahasiswa menilai bahwa pencairan dana merupakan tantangan terbesar dalam program ini, mengalahkan tantangan lainnya seperti proses seleksi, onboarding dan pelaksanaan magang itu sendiri.

3. METODE PENELITIAN

Penelitian ini dimulai dengan pengumpulan data dari platform X dalam jangka waktu Januari hingga Desember 2024 menggunakan metode crawling data python.

Metode penelitian yang digunakan dalam penelitian ini adalah Knowledge Discovery in Database yang merupakan proses penerapan teknik data mining untuk menemukan informasi berharga serta pola dalam data, dengan memanfaatkan algoritma guna mengidentifikasi pola tersebut [10].



Gambar 1. Alur penelitian

Pada gambar 1 dapat dilihat bahwa KDD sendiri memiliki 5 tahapan, yakni Data Selection, Preprocessing, Transformation, Data Mining, Evaluation.

3.1. Data Selection

Penelitian ini dimulai dengan pengumpulan data dari platform X dalam jangka waktu Januari hingga Desember 2024 menggunakan metode crawling data

python Data yang digunakan terdiri dari tweet berbahasa Indonesia yang mengandung pendapat. Pengambilan data dilakukan dengan kombinasi hashtag dan kata kunci, seperti #MSIB, #Kampus Merdeka, #MagangdanStudiIndependenBersertifikat, serta kata kunci spesifik seperti "uang saku belum cair", "BBH terlambat", "terlambat cair", "gaji tidak masuk".

3.2. Preprocessing

a. Data Cleaning

Pembersihan data akan menghilangkan beberapa data dari kumpulan data, yang memerlukan penanganan melalui berbagai metode, seperti menghapus bahasa selain Bahasa Indonesia, duplikat data, nilai yang hilang, dan penyesuaian tipe data

b. Noise Removal

Pada tahap ini, Anda akan menghilangkan karakter khusus, tanda baca, URL, dan kata-kata penghubung yang tidak diperlukan untuk analisis sentiment. Kata-kata penghubung seperti "di", "dari", "yang", dan sebagainya digunakan dalam dokumen berbahasa Indonesia.

c. Case Folding

Tahap ini dilakukan untuk mengubah seluruh huruf dalam dokumen menjadi huruf kecil dengan menggunakan bahasa pemrograman python. Proses ini bertujuan agar semua huruf sama dan memudahkan pada tahap data mining nanti. Hanya karakter dari 'a' hingga 'z' yang akan diterima, karakter selain itu akan dibuang.

d. Tokenizing

Pada tahap ini, teks akan dibagi menjadi kata-kata atau token untuk memenuhi format yang dibutuhkan oleh model. Spasi akan digunakan untuk membedakan kata-kata.

e. Stemming

Pada tahap ini, kata-kata dalam dokumen diubah ke bentuk aslinya dengan menghapus bagian tambahan seperti awalan, akhiran, dan sisipan, serta kombinasi awalan-akhiran.

3.3. Transformation

Pada tahap ini, teks yang telah melewati proses pembersihan dalam tahap pra-pemrosesan dikonversi ke dalam format numerik agar dapat diolah menggunakan metode Lexicon-Based dan Naïve Bayes Classifier. Dalam pendekatan Lexicon-Based, setiap kata dalam teks dibandingkan dengan kamus sentimen yang berisi daftar kata positif dan negative untuk menentukan skor polaritas suatu teks.

3.4. Data Mining

Pada tahap Data Mining, metode Lexicon-Based dan Naïve Bayes Classifier diterapkan untuk mengklasifikasikan sentimen mahasiswa terhadap keterlambatan pencairan uang saku MSIB berdasarkan data yang telah dikumpulkan dan diproses. Metode Lexicon-Based bekerja dengan membandingkan kata-

kata dalam teks dengan kamus sentimen, di mana setiap kata diberikan bobot berdasarkan polaritasnya, baik positif maupun negatif.

3.5. Evaluation

Tahap interpretation/evaluation menjadi tahapan terakhir dalam Knowledge Discovery in Database (KDD). Dimana pada tahap ini hasil analisis yang diperoleh dari penggabungan metode Lexicon-Based dan Naïve Bayes Classifier dievaluasi untuk memahami pola opini mahasiswa terhadap keterlambatan pencairan uang saku dalam program MSIB.

4. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan data MSIB Dalam Media Sosial X Tahun 2024. Data mentah tersebut kemudian diolah melalui tahapan Knowledge Discovery in Database (KDD), yang meliputi selection, preprocessing, transformation, data mining, dan evaluation.

4.1. Selection

Tahapan pengumpulan data menggunakan bantuan tools *Google Colab* dan bahasa pemrograman *Python*. Selain itu, API *tweet-harvest* digunakan untuk membantu proses *crawling*. Proses pengambilan data dilakukan pada beberapa tahap selama bulan Januari - Desember 2024 dengan kata kunci "Gaji," "MSIB," dan "Telat" dalam bahasa Indonesia

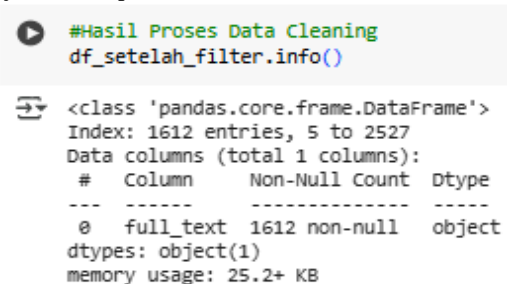
Tabel 1. Tabel Atribut

full_text
Tolong Kami Bantu Viralkan Supaya Cepat Diurus Kami Mahasiswa Magang Kampus Merdeka sudah kerja hampir 4-5 bulan tapi pencairan uang saku baru 2x dan itupun masih belum semua cair uang saku ke-2 nya [Thread & Petisi] #CairkanUangSakuMSIB#MerdekaUangSaku#KampusMerdekahttps://t.co/YIMX7XRuEY
@ditjendikti @Kemdikbud_RI @nadiemmakarim Pak please ngotak dikit udh masuk desember ini tapi uang saku yg cair baru agustus & september itupun september masih sisa ±3.000 mahasiswa belum cair kita dirantauan kesusahan ini pak tolong pak #CairkanUangSakuMSIB #KampusMerdeka https://t.co/LNDTAheWvc
Halo Diktiterus yg 28% bagaimana? sudah lewat 1 minggu semenjak yg 72% cair (18-Nov-2021) terus yg 28% kapan cairnya? 28% dari 12.900 itu masih sekitar 3600+ mahasiswa tolong janji & kewajibannya jangan diingkari terus Dikti #CairkanUangSakuMSIB #KawalUangSaku #KampusMerdeka

Tabel 1. Menjelaskan dari hasil data crawling dengan jumlah 12 atribut yang terdiri dari *created_at*, *id_str*, *full_text*, *quote_count*, *reply_count*, *retweet_count*, *favorite_count*, *lang*, *user_id_str*, *conversation_id_str*, *username*, dan *tweet_url*. Atribut yang akan digunakan hanya *full_text* yang berisi tweet dari publik.

4.2. Data Cleaning

Tahap ini dilakukan dengan menghapus bahasa selain Indonesia, data duplikat, nilai yang hilang, dan penyesuaian tipe data.



```
#Hasil Proses Data Cleaning
df_setelah_filter.info()

<class 'pandas.core.frame.DataFrame'>
Index: 1612 entries, 5 to 2527
Data columns (total 1 columns):
 #   Column      Non-Null Count  Dtype
---  ---
 0   full_text    1612 non-null   object
dtypes: object(1)
memory usage: 25.2+ KB
```

Gambar 2. Hasil data cleaning

Pada gambar 2, dijelaskan bahwa hasil yang diperoleh dari proses data cleaning yaitu sebanyak 1612 data.

4.3. Noise Removal

Selanjutnya melakukan noise removal yaitu data dilakukan pembersihan dengan tujuan untuk menghapus tanda baca, mention, hashtag, link, semua karakter yang bukan huruf (a-z atau A-Z) dan spasi serta karakter lainnya seperti '@', '/', '_', dan lainnya pada dokumen.

Tabel 2. Hasil Noise Removal

Sebelum Noise Removal	Sesudah Noise Removal
program MSIB ini kan sebenarnya buat orang orang yang ambis yang selama masa kuliahnya dipakai secara benar mengikuti organisasi cari banyak projekkan buat memperbagus portfolio dll kalau kamu nya malas-malasan selama kuliah malas bikin projek buat bagusin porto / cv terus	program MSIB ini kan sebenarnya buat orang orang yang ambil yang selama masa kuliahnya dipakai secara benar mengikuti organisasi cari banyak projekkan buat memperbagus portfolio dll kalau kamu nya malas malasan selama kuliah malas bikin projek buat bagusin porto cv terus

Pada tabel tersebut menunjukkan bahwa hasilnya adalah teks yang lebih fokus untuk meningkatkan ketepatan dan interpretabilitas

4.4. Case Folding

Tabel 3. Hasil Case Folding

Sebelum Case Folding	Setelah Case Folding
program MSIB ini kan sebenarnya buat orang orang yang ambil yang selama masa kuliahnya dipakai secara benar mengikuti organisasi cari banyak projekkan buat memperbagus portfolio dll kalau kamu nya malas malasan selama kuliah malas bikin projek buat bagusin porto cv terus	program msib ini kan sebenarnya buat orang orang yang ambil yang selama masa kuliahnya dipakai secara benar mengikuti organisasi cari banyak projekkan buat memperbagus portfolio dll kalau kamu nya malas malasan selama kuliah malas bikin projek buat bagusin porto cv terus

Kemudian dilakukan tahap Case Folding untuk menyederhanakan dan memperbarui format teks. Case folding mengonversi semua huruf dalam dokumen ke dalam format huruf kecil menggunakan fungsi lower(). Setelah itu, hanya karakter huruf dari 'a' hingga 'z' yang dipertahankan, sementara karakter selain itu dihapus.

Pada tabel 3 dapat dilihat bahwa hasilnya semua huruf dalam dokumen ke dalam format huruf kecil sehingga menjadi lebih sederhana.

4.5. Tokenizing

Tahap selanjutnya adalah Tokenizing yaitu, memecah kalimat menjadi kata atau token. Tujuannya adalah untuk mengubah data teks yang sebelumnya berbentuk kalimat panjang menjadi struktur yang lebih terorganisir sehingga lebih mudah dianalisis

Tabel 4 Hasil Tokenizing

Sebelum Tokenizing	Setelah Tokenizing
program msib ini kan sebenarnya buat orang orang yang ambil yang selama masa kuliahnya dipakai secara benar mengikuti organisasi cari banyak projekkan buat memperbagus portfolio dll kalau kamu nya malas malasan selama kuliah malas bikin projek buat bagusin porto cv terus	['program', 'msib', 'orang', 'ambis', 'kuliahnya', 'dipakai', 'mengikuti', 'ortidaknisi', 'cari', 'projekan', 'memperbagus', 'portfolio', 'nya', 'malas', 'malasan', 'kuliah', 'malas', 'bikin', 'projek', 'bagusin', 'porto', 'cv']

Pada tabel 4 dapat dilihat bahwa hasil dari tokenize menjadi lebih terorganisir.

4.6. Stemming

Tahap Stemming menggunakan library Sastrawi mengubah kata-kata dalam sebuah dokumen menjadi bentuk dasarnya dengan cara menghapus semua bagian tambahan seperti awalan, akhiran, sisipan, serta kombinasi dari awalan-akhiran khusus kata berbahasa Indonesia.

Tabel 5 Hasil Stemming

Sebelum Stemming	Setelah Stemming
['program', 'msib', 'orang', 'ambis', 'kuliahnya', 'dipakai', 'mengikuti', 'ortidaknisi', 'cari', 'projekan', 'memperbagus', 'portfolio', 'nya', 'malas', 'malasan', 'kuliah', 'malas', 'bikin', 'projek', 'bagusin', 'porto', 'cv']	['program', 'msib', 'orang', 'ambis', 'kuliah', 'pakai', 'ikut', 'ortidaknisi', 'cari', 'projekan', 'bagus', 'portfolio', 'nya', 'malas', 'malas', 'kuliah', 'malas', 'bikin', 'projek', 'bagusin', 'porto', 'cv']

4.7. Transformation

Setelah tahap preprocessing selesai dilakukan untuk membersihkan teks dari elemen-elemen yang tidak relevan, proses dilanjutkan ke tahap transformation, di mana data teks diubah menjadi

bentuk representasi yang siap dianalisis secara mendalam.

Tabel 6. Rasio Pembagian Data

Rasio Perbandingan	Data Training	Data Testing
60:40	967	645
70:30	1128	484
80:20	1289	323
90:10	1450	162

Pada tabel 6 diatas, data dibagi menjadi data latih dan data uji menggunakan fungsi `train_test_split` dari library `scikit-learn`, dengan mempertimbangkan empat variasi rasio data uji, yaitu 40%, 30%, 20%, dan 10%. Pembagian ini dilakukan untuk melihat sejauh mana ukuran data uji mempengaruhi performa model.

4.8. Data mining

Pada langkah ini, data yang telah dibagi menjadi dua bagian utama, yakni data training dan data testing dengan pembagian 60:40, 70:30, 80:20, dan 90:10, akan dilakukan implementasi klasifikasi sentiment dengan menggunakan model Naïve Bayes Classifier.

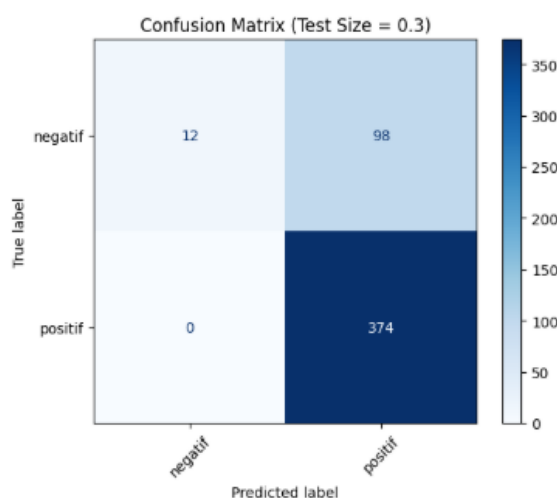
4.9. Evaluation

Dalam melakukan evaluasi hasil klasifikasi menggunakan Confusion Matrix pada berbagai pembagian ukuran uji, setiap subsampling memberikan kontribusi terhadap performa model.

Tabel 7. Hasil Evaluasi Empat Skenario

Skenario	Accuracy	Precision	Recall	F1-Score
60:40	78%	85%	53%	51%
70:30	79%	89%	55%	54%
80:20	78%	88%	52%	47%
90:10	79%	79%	55%	53%

Pada tabel 7 diatas menjelaskan bahwa dari hasil evaluasi tersebut, dapat ditentukan subsampling mana yang memberikan hasil terbaik.



Gambar 3. Hasil confusion matrix rasio 70.30

Pada gambar 3. dapat dilihat bahwa pada rasio 70:30 data yang berhasil diprediksi benar ke dalam kelas positif sebanyak 374 data, data yang berhasil diprediksi dengan benar ke dalam kelas negatif sebanyak 12 data, tidak ada kelas positif yang tidak berhasil diprediksi dengan benar, dan data kelas negatif yang tidak berhasil diprediksi dengan benar sebanyak 98 data.

5. KESIMPULAN DAN SARAN

Hasil pengolahan opini atau sentimen melalui media sosial X menggunakan algoritma Naïve Bayes Classifier menggunakan metodologi KDD menghasilkan dua kelas positif dan negatif. Hasilnya adalah 1244 label positif dan 368 label negatif. Dari perolehan sentimen tersebut, kelas sentimen positif mempunyai data yang paling banyak dan dapat diartikan meskipun ada isu keterlambatan BBH dalam program MSIB 2024, masih banyak peserta yang memberikan sentimen positif terhadap program ini, terutama terkait pengalaman magang dan semangat merdeka belajar yang diusung.

Untuk menguji tingkat akurasi analisis sentimen pada program Msib tahun 2024 menggunakan lexicon-based dan Naive Bayes Classifier, metode split data digunakan dengan empat model (90:10, 80:20, 70:30, dan 60:40). Untuk menguji hasil, matrix confusion digunakan untuk menguji skor; nilai akurasi, ketepatan, recall, dan F1-score tertinggi terdapat pada model 90:10, yang dianggap sebagai model klasifikasi terbaik. Memuat kesimpulan yang diperoleh dan saran-saran untuk penelitian selanjutnya (jika ada). Tuliskan Kesimpulan dan saran dalam 1 paragraf.

DAFTAR PUSTAKA

- [1] A. A. Azelia and H. Azzahra, "Analisis Efektivitas Implementasi Program MSIB Dalam Upaya Meningkatkan Kualitas SDM Tenaga Kerja Perguruan Tinggi Indonesia," *Inov. Makro Ekon.*, vol. 6, no. 3, pp. 183–195, 2024.
- [2] An Nisaa' Budi Sulistyaningrum, Nurulita Artanti Nirwana, Dhiya Ratri Januar, and Nela Najwa Hilalia, "Performa Kebijakan Kampus Merdeka pada Program Magang dan Studi Independen Bersertifikat," *J. Multidisiplin Madani*, vol. 2, no. 6, pp. 2771–2786, 2022, doi: 10.55927/mudima.v2i6.489.A. D. Akmal, I. Permana, H. Fajri, and Y. Yuliarti, "Opini Masyarakat Twitter terhadap Kandidat Bakal Calon Presiden Republik Indonesia Tahun 2024," *J. Manaj. dan Ilmu Adm. Publik*, vol. 4, no. 4, pp. 292–300, 2022, doi: 10.24036/jmiap.v4i4.160.
- [3] A. Rahman, D. C. Sukmajati, M. Mawar, E. Satispi, and D. Gunanto, "Implementasi Kebijakan pada Program Magang dan Studi Independen Bersertifikat di Indonesia," *SOSIOHUMANIORA J. Ilm. Ilmu Sos. Dan Hum.*, vol. 9, no. 2, pp. 266–291, 2023, doi: 10.30738/sosio.v9i2.14832.

- [4] ADAM, “ANALISIS SENTIMEN PESERTA MAGANG STUDI INDEPENDEN BERSERTIFIKAT (MSIB) KEMENDIKBUD PADA TWITTER MENGGUNAKAN NAÏVE BAYES ANALISIS SENTIMEN PESERTA MAGANG STUDI INDEPENDEN BERSERTIFIKAT (MSIB) KEMENDIKBUD PADA TWITTER MENGGUNAKAN NAÏVE BAYES,” 2023.
- [5] N. CAHYA, “Analisis Sentimen Mengenai Program MSIB Pada Media Sosial Twitter Menggunakan Algoritma Naive Bayes Classifier,” *AT-TAWASSUTH J. Ekon. Islam*, vol. VIII, no. I, pp. 1–19, 2023.
- [6] M. Hermansyah, F. Firdausi, A. Wahid, and N. Andita Prasetyo, “Twitter Sentiment Analysis for Exploring Public Opinion on the Merdeka Belajar-Kampus Merdeka (MBKM) 2023 with the Naïve Bayes Classifier Algorithm,” *Proceeding Int. Conf. Econ. Bus. Inf. Technol.*, vol. 4, pp. 852–860, 2023, doi: 10.31967/prmandala.v4i0.841C.
- F. Hasri and D. Alita, “Penerapan Metode Naïve Bayes Classifier Dan Support Vector Machine Pada Analisis Sentimen Terhadap Dampak Virus Corona Di Twitter,” *J. Inform. dan Rekayasa Perangkat Lunak*, vol. 3, no. 2, pp. 145–160, 2022, doi: 10.33365/jatika.v3i2.2026.
- [7] A. Rosadi *et al.*, “Analisis Sentimen Berdasarkan Opini Pengguna pada Media Twitter Terhadap BPJS Menggunakan Metode Lexicon Based dan Naïve Bayes Classifier Twitter Text Mining,” vol. 20, pp. 39–52, 2021.
- [8] E. S. Muhammad Bagus Dikal Putra, “METODE LEXICON BASED UNTUK ANALISIS SENTIMEN PENGGUNA TWITTER TERHADAP KINERJA ISP,” vol. 8, no. 6, pp. 12079–12087, 2024.
- [9] M. Muttaqin *et al.*, *Pengenalan Data Mining*, 1st ed. Yayasan Kita Menulis, 2023.
- [10] M. Al Khadafi, Kurnia Paranitha Kartika, and Filda Febrinita, “Penerapan Metode Naïve Bayes Classifier Dan Lexicon Based Untuk Analisis Sentimen Cyberbullying Pada Bpjs,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 6, no. 2, pp. 725–733, 2022, doi: 10.36040/jati.v6i2.5633.