

ANALISIS SENTIMEN TERHADAP PROVIDER XYZ DI TWITTER MENGUNAKAN ALGORITMA NAÏVE BAYES CLASSIFIER

Rena Muliani, Arip Solehudin, Asep Jamaludin

Program Studi Informatika S1, Fakultas Ilmu Komputer, Universitas Singaperbangsa Karawang
Jl. HS. RonggoiWaluyo, Telukjambe Timur, Karawang, Indonesia
1910631170119@student.unsika.ac.id

ABSTRAK

Kebutuhan akan jaringan internet yang semakin meningkat membuat para penyedia layanan operator seluler berlomba-lomba untuk menyediakan layanan terbaik. Telkomsel sebagai penyedia jasa seluler dengan pelanggan terbanyak di Indonesia meluncurkan XYZ, operator digital pertama yang dikembangkan untuk memenuhi kebutuhan Gen Z di Indonesia dengan layanan berbasis digital secara menyeluruh. Sebagai operator seluler XYZ kerap kali mendapatkan reaksi dari pengguna mengenai layanan yang diberikan melalui media sosial *Twitter*. Analisis sentimen menjadi sebuah solusi dalam melakukan pengolahan terhadap opini pengguna dan mengklasifikasikannya menjadi sebuah sentimen positif dan negatif dengan *text mining*. Tujuan dari penelitian ini adalah untuk mengklasifikasikan sentimen pengguna layanan *provider* XYZ menggunakan algoritma *Naïve Bayes Classifier*. Metode yang digunakan dalam penelitian ini adalah metode KDD (*Knowledge Discovery in Databases*) yang terdiri dari tahap *Data Selection*, *Preprocessing*, *Transformation*, *Data mining*, dan *Evaluation*. Hasil penelitian menunjukkan bahwa performa dari algoritma *Naïve Bayes Classifier* dalam mengklasifikasikan sentimen pengguna *provider* XYZ memiliki nilai akurasi sebesar 92%. Berdasarkan pelabelan manual, sentimen pada opini pengguna layanan *provider* XYZ lebih banyak mengandung sentimen negatif. Terlihat bahwa terdapat banyak pengguna layanan *provider* XYZ yang mengeluhkan harga, kecepatan internet, hingga ketersediaan jaringan internet. Dalam hal ini, XYZ dapat menjadikan penelitian ini sebagai acuan untuk melakukan perbaikan terhadap layanan yang diberikan.

Kata kunci: Analisis Sentimen, Confusion Matrix, KDD (*Knowledge Discovery in Database*), *Naïve Bayes Classifier*

1. PENDAHULUAN

Telkomsel resmi merilis XYZ, operator digital pertama yang dikembangkan untuk memenuhi kebutuhan Gen Z di Indonesia, XYZ menyediakan layanan berbasis digital secara menyeluruh dengan *tagline* “Semuanya Semaunya” untuk menggambarkan kebebasan dalam mengatur layanan sesuai kebutuhan dan keinginan pengguna XYZ [1]. Seluruh layanan pada kartu ini sudah berbasis digital dan dapat diakses melalui aplikasi XYZ sehingga sangat mudah dan praktis untuk dinikmati oleh penggunaannya. Jaringan pada kartu ini didukung oleh jaringan 4G hingga 5G dari Telkomsel membuat penggunaannya dapat merasakan koneksi internet yang stabil dengan cakupan wilayah yang luas [2].

Berdasarkan pada situs resmi *Provider* XYZ, per tanggal 22 Agustus 2022, ada aturan batas penggunaan wajar atau FUP (*Fair Usage Policy*) sebesar 100 GB per paket Mbps aktif. Jika sudah mencapai batas penggunaan, pengguna tetap bisa menggunakan internet tetapi hanya dengan batas kecepatan maksimum 128 Kbps [3]. Selain itu, pada bulan April 2023 Telkomsel sudah tidak lagi bekerja sama dengan *Disney+ Hotstar*, berakhirnya masa promo paket kuota dengan bonus berlangganan *Disney+ Hotstar* pada bulan April 2023. Hal-hal tersebut menuai cukup banyak keluhan dari pengguna yang dituangkan pada media sosial salah satunya pada *Twitter*.

Twitter berisi banyak sekali kiriman teks dan akan terus bertambah setiap harinya dan pengguna media sosial ini juga sangat beragam dari pengguna biasa hingga selebriti, politisi, perwakilan perusahaan, bahkan presiden dari suatu negara. Oleh karena itu, sangat memungkinkan untuk mengumpulkan *tweet* dari pengguna dengan berbagai macam kelompok sosial dan minat yang berbeda-beda [4]. Apabila dibandingkan dengan data *platform* digital lainnya seperti Facebook, Instagram, Snapchat, dan lain-lain, data pada *Twitter* lebih mudah diakses dan berisi sumber daya yang berharga untuk melakukan penelitian akademis [5].

Twitter kerap kali dijadikan tempat untuk mengeluarkan opini mengenai berbagai hal oleh penggunaannya, termasuk opini mengenai performa XYZ sebagai operator seluler, terlebih ketika kebijakan mengenai FUP diberlakukan serta berakhirnya masa promo paket kuota dengan bonus berlangganan *Disney+ Hotstar* berakhir. XYZ dapat terus meningkatkan kualitas dan performa jika mengetahui opini dari penggunaannya. Meski begitu belum dapat diketahui apakah opini dari pengguna terkait XYZ cenderung positif atau negatif. Hal tersebut dapat diatasi dengan melakukan pengolahan terhadap opini dengan *text mining* dan melakukan analisis sentimen.

Analisis sentimen merupakan bagian dari *text mining*, dimana data yang akan diolah adalah data berupa teks [20]. Penelitian mengenai analisis

sentimen sebelumnya telah dilakukan oleh Niqam Haqqizar dan Tiqa Nur L. (2019) mengenai analisis terhadap layanan *Provider* Telkomsel menggunakan algoritma *Naïve Bayes* pada media sosial *Twitter* dengan data yang diambil sebanyak 151 data. Hasil dari pengolahan data tersebut menunjukkan bahwa banyak masyarakat yang memberikan sentimen negatif terhadap layanan Telkomsel, dengan 51 *tweet* sentimen negatif, 49 sentimen positif, dan 51 sentimen netral yang merupakan opini tidak bersentimen. Tingkat akurasi dalam penentuan kategori tersebut sebesar 70,21%, serta *micro average* sebesar 70,20% dalam penentuan sentimen memiliki tingkat *Precision* 70,11% dan *Recall* 70,33%.

Berdasarkan pembahasan di atas maka akan dilakukan sebuah penelitian mengenai bagaimana analisis sentimen terhadap layanan XYZ dengan cara mengklasifikasikan opini pengguna XYZ serta menganalisis dan mengevaluasi hasil analisis sentimen pengguna XYZ dengan menggunakan algoritma *Naïve Bayes*.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data mining adalah serangkaian tahapan yang bertujuan untuk menemukan informasi bernilai tambah yang sebelumnya tidak dapat diketahui secara manual dari sebuah *database*. Informasi berharga tersebut ditemukan dengan memeriksa dan menganalisis pola yang menarik atau penting dari kumpulan data yang disimpan dalam *database*. *Data mining* sering dimanfaatkan untuk mencari informasi dari basis data berukuran besar, juga dikenal sebagai *Knowledge Discovery Databases (KDD)* [6].

2.2. Text Mining

Text mining adalah suatu teknik yang berfungsi untuk memproses *information extraction*, *information retrieval*, klasifikasi, dan *clustering*. Tidak seperti *data mining*, *text mining* memanfaatkan pola bahasa alami yang tidak terstruktur, sedangkan *data mining* menggunakan pola yang diperoleh dari basis data yang memiliki struktur yang sudah ditentukan [7].

2.3. Analisis Sentimen

Metode analisis sentimen digunakan untuk mengevaluasi penilaian masyarakat, opini, perilaku, dan emosi terhadap berbagai entitas seperti produk, individu, organisasi, layanan publik, kejadian, topik, dan isu. Fokus utama analisis sentimen adalah pada pesan positif atau negatif yang terkandung dalam suatu opini [8].

2.4. Klasifikasi

Klasifikasi merupakan bagian dari pembelajaran *machine learning* yang mempelajari data yang telah diklasifikasikan. Ada sejumlah algoritma yang dapat digunakan pada klasifikasi sentimen seperti *Naïve Bayes*, *K-Nearest Neighbor (KNN)*, *Support Vector*

Machine, *Logistic Regression*, *Artificial Neural Network*, dan sebagainya [8].

2.5. Naïve Bayes Classifier

Algoritma klasifikasi *Naïve Bayes Classifier* (NBC) merujuk pada teorema Bayes di bidang statistika. *Naïve Bayes Classifier* dikategorikan sebagai algoritma klasifikasi linier yang efisien [8]. Dalam penggunaannya, algoritma ini lebih sesuai untuk menangani *dataset* yang memiliki tipe nominal [16].

Tujuan dari pembelajaran menggunakan *Naïve Bayes Classifier* adalah untuk memprediksi probabilitas masing-masing kategori yang ada di setiap ciri dokumen yang diuji. Rumus untuk metode ini telah disederhanakan dalam persamaan berikut.

$$P(c_j|w_i) = \frac{P(w_i|c_j)P(c_j)}{P(w_i)} \quad (1)$$

Keterangan

$P(c_j|w_i)$: Probabilitas kategori j bila kata i muncul

$P(w_i|c_j)$: Probabilitas kata i termasuk ke dalam kategori j

$P(c_j)$: Probabilitas munculnya kategori j

$P(w_i)$: Probabilitas munculnya sebuah kata

2.6. Confusion Matrix

Confusion Matrix adalah sebuah metode yang biasa digunakan untuk mengukur seberapa baik *Classifier* bisa mengenali tupel dari kelas yang berbeda. Untuk mempermudah dalam pemahaman penggunaan *Confusion Matrix* dibagi menjadi *matrix* 2×2 sebagai berikut [9].

		Predicted Class		Total
		True	False	
Actual Class	True	TP	FN	P
	False	FP	TN	N
Total		P'	N'	P + N

Sumber: Han et al. (2014)

Gambar 1. Table confusion matrix

2.7. Area Under Curve (AUC)

AUC adalah perhitungan sebagian dari satuan luas persegi yang didapatkan dari grafik ROC (*Receiver Operating Characteristic*) curve yang nilainya akan selalu berada di antara 0,0 dan 1,0. ROC curve didapatkan melalui grafik yang dimana pada sumbu-x (Horizontal) menggunakan perhitungan FP (*False Positive*) Rate dan pada sumbu-y (Vertikal) menggunakan perhitungan dari TP (*True Positive*) Rate dengan rumus perhitungannya sebagai berikut [10].

Tabel 1. Kualitas *classifier* dengan AUC

Nilai AUC	Kualitas
0.90 – 1.00	<i>Excellent classification</i>
0.80 – 0.90	<i>Good classification</i>
0.70 – 0.80	<i>Fair classification</i>
0.60 – 0.70	<i>Poor classification</i>
0.50 – 0.60	<i>Failure</i>

Sumber: Gorunescu (2010)

2.8. Term Frequency – Inverse Document Frequency (TF-IDF)

Term Frequency-Inverse Document Frequency (TF-IDF) adalah sebuah skema pembobotan kata yang umum digunakan. Metode TF-IDF terkenal karena kemudahan penggunaannya, efisiensi, dan memberikan hasil yang tepat atau akurat. Metode ini digunakan untuk perhitungan nilai *Term Frequency* (TF) dan *Inverse Document Frequency* (IDF) untuk setiap kata dalam sebuah dokumen [11]. TF (*Term Frequency*) adalah rasio dari jumlah suatu kata pada kalimat dibandingkan dengan panjang dari kalimat tersebut. IDF (*Inverse Document Frequency*) dari setiap kata adalah log dari rasio berdasarkan total dokumen dengan jumlah dari dokumen tertentu yang dimana kata tersebut ada. TF-IDF adalah gabungan dari TF dan IDF yang dimana keduanya bisa menutupi kekurangan untuk membuat prediksi menjadi relevan [17].

Nilai *Term Frequency* (TF) didapatkan dari nilai frekuensi munculnya fitur t pada dokumen d [18].

$$TF_{t,d} = F(t, d) \quad (2)$$

Nilai *Inverse Document Frequency* (IDF) dapat dihitung dari logaritma banyaknya dokumen N dibagi banyaknya dokumen DF_t yang memiliki fitur t .

$$IDF_t = \log \frac{N}{DF_t} \quad (3)$$

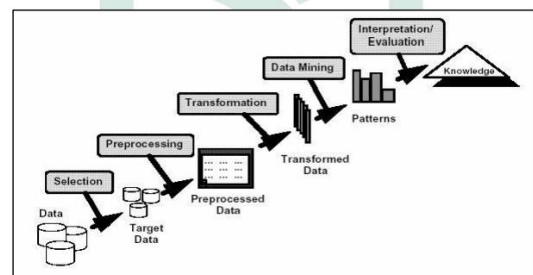
Nilai TF-IDF (W_t) didapatkan dengan mengalikan nilai TF (persamaan 2.21) dengan nilai IDF (persamaan 2.22) [18].

$$W_t = TF_{t,d} \times IDF_t \quad (4)$$

Hasil dari tahap pembobotan ini selanjutnya akan digunakan pada tahap klasifikasi dokumen.

2.9. Knowledge Discovery in Database (KDD)

Knowledge Discovery in Database (KDD) adalah proses yang kompleks untuk menemukan dan mendeteksi pola pada data yang dapat dimengerti, baru, valid, dan bermanfaat. Proses *Knowledge Discovery* meliputi hasil dari proses *data mining*, dimana hasilnya kemudian diubah menjadi informasi yang mudah dimengerti secara akurat [12]. Secara umum, proses *Knowledge Discovery in Database* mencakup beberapa tahap, yaitu [19]:

Gambar 2. Proses *Knowledge Discovery in Database*

Berdasarkan Gambar 2, proses *Knowledge Discovery in Database* memiliki tahapan-tahapan sebagai berikut:

1. **Data Selection**
Proses seleksi data dilakukan sebelum memulai tahap KDD. Data yang terpilih akan digunakan dalam proses *data mining* dan ditempatkan dalam sebuah file yang berbeda dari data operasional.
2. **Pre-processing**
Proses preprocessing atau Cleaning merupakan proses membersihkan data yang mencakup penghapusan data duplikat, pengecekan data yang tidak konsisten, dan mengoreksi kesalahan pada data. Selain itu, juga dilakukan proses *enrichment* yaitu proses pengayaan data yang tersedia dengan data atau informasi tambahan yang berkaitan dan diperlukan untuk *Knowledge Discovery in Database*.
3. **Transformation**
Proses memilih data yang cocok untuk proses *data mining*. Proses *transformation* merupakan proses yang kreatif dan sangat tergantung pada jenis informasi atau pola yang ingin ditemukan dalam *database*.
4. **Data mining**
Proses pencarian pola atau informasi menarik yang terdapat dalam data yang terpilih melalui metode atau teknik tertentu. Metode, algoritma, atau teknik dalam *data mining* cukup beragam. Pemilihan algoritma atau metode yang tepat sangat bergantung pada tujuan dan keseluruhan proses *Knowledge Discovery in Database*.
5. **Interpretation/Evaluation**
Tahap *interpretation* menampilkan pola informasi hasil dari proses *data mining* dalam bentuk yang mudah dipahami. Pada tahap ini dilakukan pengecekan apakah pola atau informasi yang ditemukan sesuai atau bertentangan dengan fakta atau hipotesis yang ada sebelumnya.

2.10. Python

Python adalah salah satu bahasa pemrograman yang memiliki tujuan umum terbaik yang bisa digunakan di berbagai sistem operasi modern. *Python* juga merupakan bahasa pemrograman dengan tipe *interpreted language* yang dimana *code* tidak akan diubah menjadi sebuah *code* yang dapat dipahami oleh komputer sebelum dijalankan, melainkan perubahan tersebut terjadi ketika *runtime*. *Python* diciptakan oleh

Guido Van Rossum. Aplikasi yang dapat dikembangkan dari aplikasi ini beragam, salah satunya di bidang *Scientific* dan *numeric* yang memiliki berbagai *library* Scipy, Pandas, dan Scikit-learn [13].

2.11. Twitter

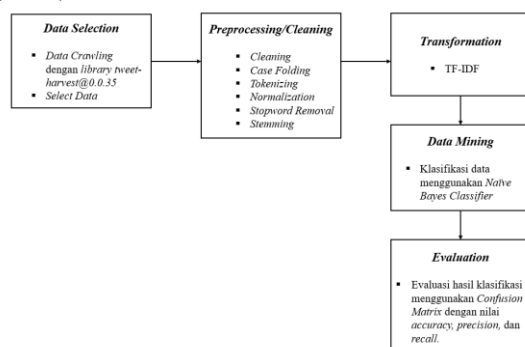
Twitter merupakan suatu *platform* media sosial bertipe *micro-blogging* atau blog berukuran kecil yang didirikan oleh Jack Dorsey dan dirilis pada bulan Juli tahun 2006. Twitter merupakan salah satu media sosial yang menarik perhatian bagi pendidik untuk berbagai penggunaan profesional ataupun untuk melakukan suatu penelitian [14].

2.12. Google Colaboratory

Google Colaboratory atau biasa disebut Google Colab merupakan sebuah IDE untuk pemrograman Python dengan pemrosesan yang dilakukan oleh server Google yang memiliki perangkat keras dengan performa tinggi. Dari sisi perangkat lunak, Google Colab telah menyediakan hampir sebagian besar *library* yang dibutuhkan. Dengan jaminan kemampuan servernya yang stabil hampir keseluruhan pemrosesan tidak menemukan kendala dengan Google Colab selama koneksi jaringan internet lancar [15].

3. METODE PENELITIAN

Penelitian ini dilakukan dengan menggunakan metodologi *Knowledge Discovery in Database* (KDD).



Gambar 3. Alur penelitian

Berdasarkan Gambar 3, metodologi penelitian yang akan digunakan dalam penelitian ini adalah metodologi KDD yang terdiri dari 5 tahapan. Secara umum, proses *Knowledge Discovery in Database* mencakup beberapa tahap, yaitu *data selection*, *preprocessing*, *transformation*, *data mining*, dan *evaluation* [19].

3.1. Data Selection

Tahap *data selection* akan dilakukan proses mengumpulkan dan memilih data yang akan digunakan untuk penelitian, tahapan *data selection* adalah sebagai berikut:

3.2. Data Crawling

Proses pengumpulan data dalam penelitian ini akan dilakukan dengan menggunakan *library tweet-harvest@0.0.35* dengan menggunakan NodeJS. *Library tweet-harvest@0.0.35* digunakan dengan menggunakan *auth token* dari Twitter. Data yang akan digunakan merupakan data *tweet* dari pengguna layanan XYZ yang diambil dari media sosial Twitter dan dikhususkan pada rentang tanggal 1 Februari 2023 sampai dengan 1 Mei 2023 dengan menggunakan kata kunci yang sudah dijelaskan pada sub bab 3.1.

3.3. Select Data

Tahap *select data* dilakukan untuk memilih data yang relevan serta atribut yang digunakan untuk masuk ke dalam tahap *preprocessing* dengan cara menghapus data atau atribut yang tidak digunakan. Data akan dibagi menjadi 2 kategori yaitu Positif dan Negatif. Kemudian akan dilakukan pelabelan data yang akan dibantu melalui validasi oleh ahli bahasa berdasarkan makna kalimat untuk meminimalisir kesalahan pada pelabelan.

3.4. Preprocessing/Cleaning

Pembersihan data dilakukan karena banyak data yang menggunakan bahasa yang tidak sesuai maupun terdapat beberapa atribut data yang tidak diperlukan dalam proses klasifikasi pada penelitian ini. *Preprocessing* melalui beberapa tahap yaitu *Cleaning*, *case folding*, *Tokenizing*, *Normalization*, *stopword removal*, dan *Stemming*.

3.5. Transformation

Tahap *transformation* ini dilakukan seleksi fitur menggunakan TF-IDF (*Term Frequency Inverse Document Frequency*). Seleksi fitur dengan TF-IDF ini dilakukan dengan memberikan bobot terhadap kemunculan kata pada sebuah data. Proses ini bertujuan untuk mengurangi jumlah variabel dan data yang tidak terlalu diperlukan.

3.6. Data Mining

Tahap *data mining* merupakan tahap dilakukannya proses klasifikasi *tweet* pengguna Provider XYZ menggunakan algoritma *Naïve Bayes Classifier*. Tahap klasifikasi dilakukan berdasarkan data yang telah divalidasi oleh ahli kebahasaan berdasarkan konteks kalimat. Data yang ada kemudian dipecah menjadi dua bagian, yaitu data *training* dan data *testing*. Pada penelitian ini digunakan 4 skenario penggunaan data *training* dan data *testing*. 4 skenario tersebut sebagai berikut:

Skenario pertama : 90% data *training* dan 10% data *testing*

Skenario kedua : 80% data *training* dan 20% data *testing*

Skenario ketiga : 70% data *training* dan 30% data *testing*

Skenario keempat: 60% data *training* dan 40% data *testing*

Setelah data dibagi menjadi data *training* dan data *testing*, tahap selanjutnya adalah menerapkan algoritma *Naïve Bayes Classifier* dalam proses *data mining*.

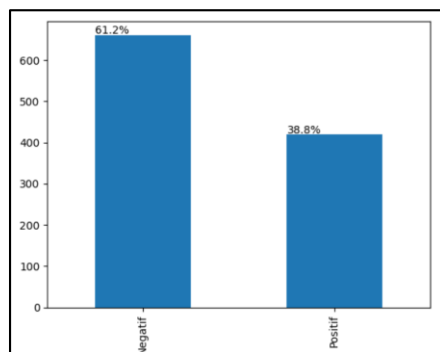
3.7. Evaluation

Tahap *evaluation* merupakan tahap terakhir dalam metode *Knowledge Discovery in Database* (KDD) yang mencakup nilai hasil kesimpulan dari pengolahan *data mining*. Tahap evaluasi dilakukan untuk mengukur validitas hasil dari klasifikasi dengan menghitung nilai akurasi menggunakan *Confusion Matrix*. Dalam *Confusion Matrix* parameter yang dihitung adalah nilai *accuracy*, *Precision*, *Recall*, dan *F1-Score*.

4. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan algoritma *Naïve Bayes Classifier* dalam membuat model analisis sentimen terhadap opini pengguna *Provider* digital XYZ dengan menggunakan metode *Knowledge Discovery in Databases* (KDD) yang diimplementasikan melalui lima tahap yaitu *Data Selection*, *Preprocessing*, *Transformation*, *Data mining*, dan *Evaluation*. Data yang digunakan dalam penelitian ini total 1080 data *tweet* yang diperoleh dari *Twitter* dengan rentang waktu 1 Februari 2023 hingga 1 Mei 2023. Penelitian dimulai dengan proses pengumpulan data menggunakan *package tweet-harvest@0.0.35* pada *NodeJs*. Data yang dikumpulkan merupakan data opini pengguna mengenai layanan *Provider XYZ* pada *platform Twitter*. Kata kunci yang digunakan dalam pengumpulan data adalah “@byu_id” dan “#byu.u”.

Setelah proses pengumpulan data selesai, dilakukan pemilihan data dengan memilih atribut “text” yang berisi *tweet* atau opini yang diteliti, kemudian melakukan pelabelan data secara manual dengan bantuan validasi dari seorang ahli bahasa. Data opini pengguna XYZ dibagi ke dalam dua kelas yaitu positif dan negatif.



Gambar 4. Presentase jumlah label sentimen

Gambar 4 menunjukkan persentase dari jumlah sentimen positif dan sentimen negatif yang terdapat pada *tweet* terkait layanan *Provider XYZ*. Pada

penelitian ini kelas positif meliputi pujian, ungkapan terima kasih, ungkapan senang, dan antusiasme positif, sedangkan kelas negatif meliputi kritik, umpatan, keluhan, dan sarkasme. Pelabelan yang dilakukan menghasilkan persentase data negatif yang lebih besar dibandingkan dengan persentase data positif.

Hasil dari pelabelan data tersebut kemudian dilanjutkan ke tahap *Preprocessing* untuk membersihkan data. *Preprocessing* melibatkan enam tahap yang terdiri dari *Cleaning*, *case folding*, *Tokenizing*, *Normalization*, *stopword removal*, dan *Stemming*. Tahap tersebut dilakukan bertujuan untuk memastikan data siap untuk dianalisis. Hasil dari tahap *Preprocessing* kemudian dilanjutkan untuk proses transformasi data menggunakan metode TF-IDF (*Term frequency – Inverse Document Frequency*) untuk memberikan bobot pada setiap kata dalam *dataset* dan menghasilkan representasi numerik untuk setiap kata. Selanjutnya untuk tahap *Data mining* dilakukan dengan menggunakan algoritma *Naïve Bayes Classifier*.

Data kemudian dibagi menjadi data *training* dan data *testing* dengan empat skenario yaitu 90:10, 80:20, 70:30, dan 60:40. Kemudian dari masing-masing skenario dilakukan evaluasi performa dengan menggunakan *Confusion Matrix* untuk mengetahui nilai *accuracy*, *Precision*, *Recall*, dan *F1-Score*.

Tabel 1 menunjukkan skenario pengujian data *testing* dan data *training*, data tersebut diambil dari total 1080 data yang telah diproses sebelumnya dengan pembagian sebagai berikut.

Tabel 2. Skenario pengujian data *training* dan *testing*

Persentase Data		Jumlah Data	
Training	Testing	Training	Testing
90%	10%	972	108
80%	20%	864	216
70%	30%	756	324
60%	40%	648	432
Total		1080	

Berikut hasil nilai *accuracy*, *precision*, *recall*, *f1-score*, dan nilai AUC ROC yang didapatkan dari setiap skenario.

4.1. Skenario 1 (90:10)

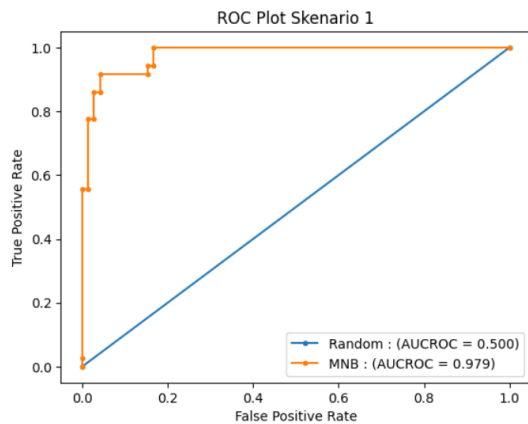
Berikut dapat dilihat pada Gambar 5 hasil evaluasi dari skenario 1.

	precision	recall	f1-score	support
Negatif	0.90	0.97	0.93	72
Positif	0.93	0.78	0.85	36
accuracy			0.91	108
macro avg	0.92	0.88	0.89	108
weighted avg	0.91	0.91	0.91	108

Sentimen Positif : 36 review atau 33.3333 %
 Sentimen Negatif : 72 review atau 66.6667 %
 Negatif 661
 Positif 419

Gambar 5. Hasil evaluasi skenario 1 (90:10)

Berdasarkan Gambar 5, skenario pertama yang menggunakan data *training* 90% dan data *testing* 10% menghasilkan nilai *accuracy* sebesar 91%, *Precision* 93%, *Recall* 78%, dan *F1-Score* 85%.



Gambar 6. Nilai AUC ROC skenario 1 (90:10)

Berdasarkan Gambar 6 dapat diketahui bahwa nilai AUC ROC untuk skenario pertama dengan pembagian data *training* 90% dan data *testing* 10% yaitu sebesar 0.979.

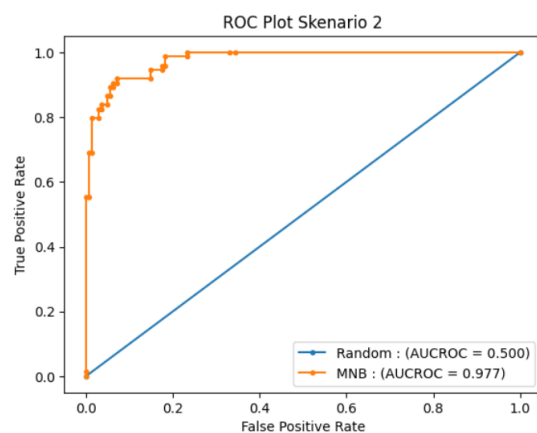
4.2. Skenario 2 (80:20)

Berikut dapat dilihat pada Gambar 7 hasil evaluasi dari skenario 2.

	precision	recall	f1-score	support
Negatif	0.90	0.99	0.94	142
Positif	0.97	0.80	0.87	74
accuracy			0.92	216
macro avg	0.94	0.89	0.91	216
weighted avg	0.93	0.92	0.92	216

Gambar 7. Hasil evaluasi skenario 2 (80:20)

Berdasarkan Gambar 7, skenario kedua yang menggunakan data *training* 80% dan data *testing* 20% menghasilkan nilai *accuracy* sebesar 92%, *Precision* 97%, *Recall* 80%, dan *F1-Score* 87%.



Gambar 8. Nilai AUC ROC skenario 2 (80:20)

Berdasarkan Gambar 8 dapat diketahui bahwa nilai AUC ROC untuk skenario kedua dengan

pembagian data *training* 80% dan data *testing* 20% yaitu sebesar 0.977.

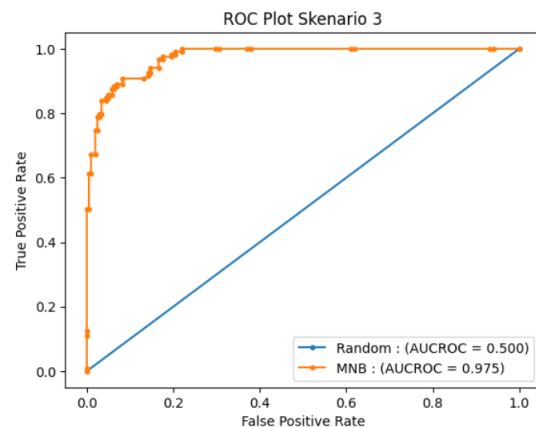
4.3. Skenario 3 (70:30)

Berikut dapat dilihat pada Gambar 9 hasil evaluasi dari skenario 3.

	precision	recall	f1-score	support
Negatif	0.88	0.98	0.93	205
Positif	0.95	0.77	0.85	119
accuracy			0.90	324
macro avg	0.91	0.87	0.89	324
weighted avg	0.91	0.90	0.90	324

Gambar 9. Hasil evaluasi skenario 3 (70:30)

Berdasarkan Gambar 9, skenario ketiga yang menggunakan data *training* 70% dan data *testing* 30% menghasilkan nilai *accuracy* sebesar 90%, *Precision* 95%, *Recall* 77%, dan *F1-Score* 85%.



Gambar 10. Nilai AUC ROC skenario 3 (70:30)

Berdasarkan Gambar 10 dapat diketahui bahwa nilai AUC ROC untuk skenario ketiga dengan pembagian data *training* 70% dan data *testing* 30% yaitu sebesar 0.975.

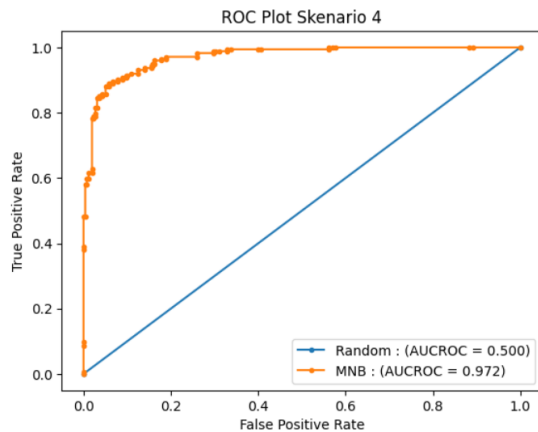
4.4. Skenario 4 (60:40)

Berikut dapat dilihat pada Gambar 11 hasil evaluasi dari skenario 4.

	precision	recall	f1-score	support
Negatif	0.85	0.98	0.91	258
Positif	0.96	0.74	0.84	174
accuracy			0.88	432
macro avg	0.91	0.86	0.87	432
weighted avg	0.89	0.88	0.88	432

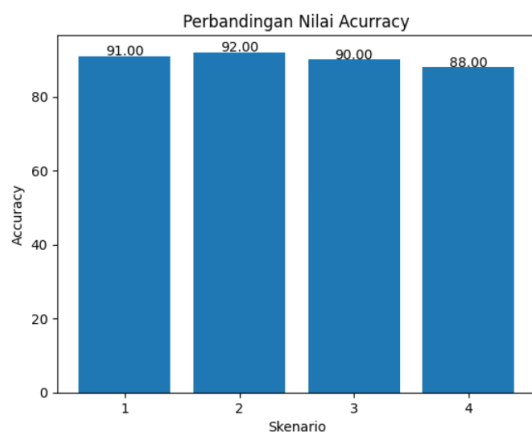
Gambar 11. Hasil evaluasi skenario 4 (60:40)

Berdasarkan Gambar 11, skenario keempat yang menggunakan data *training* 60% dan data *testing* 40% menghasilkan nilai *accuracy* sebesar 88%, *Precision* 96%, *Recall* 74%, dan *F1-Score* 84%.



Gambar 12. Nilai AUC ROC skenario 4 (60:40)

Berdasarkan Gambar 12 dapat diketahui bahwa nilai AUC ROC untuk skenario pertama dengan pembagian data *training* 60% dan data *testing* 40% yaitu sebesar 0.972.



Gambar 13. Perbandingan akurasi setiap skenario

Berdasarkan Gambar 13, hasil evaluasi nilai akurasi pada setiap skenario yang ditunjukkan pada Gambar 4.20 menunjukkan bahwa pada skenario pertama (90:10) dengan data *training* 90% dan data *testing* 10% tingkat akurasi yang dihasilkan adalah 91%. Pada skenario kedua (80:20) dengan data *training* 80% dan data *testing* 20% tingkat akurasi yang dihasilkan adalah 92%. Pada skenario ketiga (70:30) dengan data *training* 70% dan data *testing* 30% tingkat akurasi yang dihasilkan adalah 90%. Sedangkan untuk skenario keempat (60:40) dengan data *training* 60% dan data *testing* 40% tingkat akurasi yang dihasilkan adalah 88%.

Berdasarkan perbandingan hasil pengujian evaluasi, skenario terbaik untuk model klasifikasi dengan *Naïve Bayes Classifier* adalah pada skenario kedua dengan 80% data *training* dan 20% data *testing* yang memiliki nilai *accuracy* sebesar 92%, *Precision* 97%, *Recall* 80%, dan *F1-Score* 87%. Model ini dinilai baik karena mempunyai nilai evaluasi yang persentasenya sudah sangat baik.

Tabel 3. Tabel Nilai AUC ROC

Skenario	Nilai AUC ROC
1	0,979
2	0,977
3	0,975
4	0,972

Tabel 3 menunjukkan nilai AUC ROC pada setiap skenario. Berdasarkan tabel tersebut dapat dilihat bahwa skenario pertama (pembagian data *training* 90% dan data *testing* 10%) memiliki nilai AUC ROC tertinggi yaitu 0.979. Nilai AUC ROC untuk setiap skenario termasuk dalam predikat terbaik yakni *excellent classification* karena memiliki nilai di atas 0.9 berdasarkan pada referensi Gorunescu.

5. KESIMPULAN DAN SARAN

Analisis sentimen pengguna layanan *Provider XYZ* pada media sosial *Twitter* didominasi oleh sentimen negatif. Pada data yang terdiri dari kumpulan opini pengguna XYZ pada *Twitter* dari total 1080 data ditemukan jumlah sentimen positif sebanyak 419 dan jumlah sentimen negatif sebanyak 661 yang masing-masing diperoleh melalui pemilihan data secara manual yang dibantu oleh ahli Bahasa. Berdasarkan pelabelan data secara manual, layanan *Provider XYZ* memperoleh banyak sentimen negatif dilihat dari pengguna yang mengeluh mengenai kualitas jaringan dan harga pada media sosial *Twitter*.

Performa algoritma *Naïve Bayes Classifier* dalam mengklasifikasikan sentimen pengguna layanan *Provider XYZ* memiliki nilai akurasi sebesar 92% dan nilai AUC sebesar 0,977 yang termasuk dalam predikat terbaik yakni *excellent classification* yang diperoleh dari skenario kedua dengan pembagian data 80:20 yakni 80% data *training* dan 20% data *testing*. Pemberian label menggunakan metode manual dengan pembobotan TF-IDF.

Saran penelitian selanjutnya, peneliti dapat menggunakan Teknik SMOTE untuk menangani data yang tidak seimbang. Selain itu dapat menggunakan metode pembobotan dan seleksi fitur seperti *Chi Square*, *N-Gram*, *Particle Swarm Optimization*, dan lain sebagainya. Serta dapat menggunakan sumber *dataset* yang berbeda seperti pada media sosial Instagram, Facebook, atau media sosial lainnya.

DAFTAR PUSTAKA

- [1]. Telkomsel, "Telkomsel Luncurkan By.U, Layanan Seluler Prabayar digital End-to-end Pertama di Indonesia" *Telkomsel.com*, Oktober. 2019. <https://www.telkomsel.com/about-us/news/telkomsel-luncurkan-byu-layanan-selular-prabayar-digital-end-end-pertama-di-indonesia> (accessed Mar. 5, 2023).
- [2]. F. Rohman, "Cara Beli Kartu By.U Termurah Beserta Metode Aktivasinya" *Katadata.co.id*, Agustus. 2022. <https://katadata.co.id/intan/berita/62f0e5da62d91/cara-beli-kartu-byu-termurah-beserta-metode-aktivasinya> (accessed Mar. 5, 2023).

- [3]. By.U, "Apakah Paket Mbps by.U Ada FUP/Batas Penggunaan Wajarnya?" *byu.id*, Agustus. 2022. <https://www.byu.id/id/faq-detail/fup-paket-mbps> (accessed Ags. 4, 2023).
- [4]. D. A. Kristiyanti, A. H. Umam, R. Amin, and L. Marlinda, "Comparison of SVM Naïve Bayes Algorithm for Sentiment Analysis Toward West Java Governor Candidate Period 2018-2023 Based on Public Opinion on Twitter," *6th International Conference on Cyber and IT Service Mangement (CITSM 2018)*, Aug 7-9. 2018, doi: 10.1109/CITSM.2018.8674352
- [5]. J. Yu, and J. Munoz-Justicia, "A Bibliometric Overview of Twitter-Related Studies Indexed in Web of Science," *Future Internet*, vol. 12, no. 91, May. 2020, doi: <https://doi.org/10.3390/fi12050091>.
- [6]. Vlandari, 2017. "Data Mining Teori dan Aplikasi Rapidminer," Yogyakarta: Gava Media.
- [7]. D. A. Agustina, S. Subanti, and E. Zhukronah, "Implementasi Text Mining Pada Analisis Sentimen Pengguna Twitter Terhadap Marketplace di Indonesia Menggunakan Algoritma Support Vector Machine," *Indonesian Journal of Applied Statistics*, vol. 3, no. 2, pp. 109-122, 2020.
- [8]. A. Muzaki and A. Witanti, "Sentimen Analisis Masyarakat Di Twitter Terhadap Pilkada 2020 Ditengah Pandemi Covid-19 Dengan Metode Naïve Bayes Classifier," *Jurnal Teknik Informatika (JUTIF)*, vol. 2, no. 2, pp. 101-107, 2021.
- [9]. J. Han, M. Kamber, and J. Pei, "Data Mining: Concept and Techniques" In *Data Mining: Concepts and Techniques (Third)*. United States of America: Elsevier Inc. 2012, doi: <https://doi.org/10.1016/C2009-0-61819-5>
- [10]. F. Gorunescu, "Data Mining Concepts, Model, and Techniques(J. Kacprzyk & L. C. Jain, Eds.)," Chennai, India: Springer. 2010, Accessed: Apr. 28, 2023. doi: <https://doi.org/10.1007/978-3-642-19721-5>
- [11]. Muljono, D. P. Artanti, A. Syukur, A. Prihandono, D. R. I. M. Setiadi, "Analisis Sentimen Untuk Penilaian Pelayanan Situs Belanja Online Menggunakan Algoritma Naïve Bayes," *Konferensi Nasional Sistem Informasi 2018*, pp. 165-170, 2018.
- [12]. A. A. Fajrin and A. Maulana, "Penerapan Data Mining Untuk Analisis Pola Pembelian Konsumen Dengan Algoritma FP-Growth Pada Data Transaksi Penjualan Spare Part Motor," *Kumpulan Jurnal Ilmu Komputer (KLIK)*, vol. 05, no. 01, 2018.
- [13]. Python Software Foundation, "Applications for Python" *python.org*, 2023. <https://www.python.org/about/apps/#scientific-and-numeric> (accessed Mar. 5, 2023).
- [14]. J. Carpenter, T. Tani, S. Morison, and J. Keane, "Exploring The Landscape of Educator Professional Activity on Twitter: an Analysis of 16 Education-Related Twitter Hashtags," *Professional Development in Education*, vol. 00, no. 00, pp. 1-22, 2020.
- [15]. R. T. Handayanto and H. Herlaw, "Prediksi Kelas Jamak Dengan Deep Learning Berbasis Graphics Processing Unit," *Jurnal Kajian Ilmiah*, vol. 20, no. 1, pp. 67-76, 2020.
- [16]. J. Suntoro, 2019. "Data Mining: Algoritma Dan Implementasi Dengan Pemrograman PHP," Jakarta: PT Elex Media Komputindo
- [17]. A. Kulkarni and A. Shivananda, 2019. "Natural language processing recipes," In C. S. John, M. Moodie, & S. Vishwakarma (Eds.), *Natural Language Processing Recipes*. New York: Springer Science+, doi: <https://doi.org/10.1007/978-1-4842-4267-4>
- [18]. N. A. Verdikha, T. B. Adji, and A. E. Permanasari, "Komparasi Metode Oversampling Untuk Klasifikasi Teks Ujaran Kebencian," *Seminar Nasional Teknologi Informasi dan Multimedia 2018*, Feb 10, 2020.
- [19]. F. A. Hermawati, 2013. "Data Mining," Yogyakarta: Andi Offset.
- [20]. N. Haqqizar and T. N. Larasyanti, "Analisis Sentimen Terhadap Layanan Provider Telekomunikasi Telkomsel Di Twitter Dengan Metode Naïve Bayes," *Prosiding TAU SNAR-TEK 2019 Seminar Nasional Rekayasa dan Teknologi*, pp. 30-33, 2019