

ANALISIS SENTIMEN PENGGUNA SLIK OJK MENGGUNAKAN NAÏVE BAYES DENGAN OPTIMASI INFORMATION GAIN DAN SMOTE

Tatan Darusalam, Syariful Alam, Mutiara Andayani Komara

Teknik Informatika, Sekolah Tinggi Teknologi Wastukencana Purwakarta
Jl.Cikopak No.53, Kec Purwakarta, Purwakarta, Jawa Barat, Indonesia-41151
tatandarusalam33@wastukencana.ac.id

ABSTRAK

Bank Indonesia Checking (BI Checking) atau Sistem Informasi Debitur (SID) sangat penting untuk verifikasi data nasabah dan mitigasi risiko bagi penyedia layanan fintech. Pada 27 April 2017, Otoritas Jasa Keuangan (OJK) memperkenalkan Sistem Layanan Informasi Keuangan (SLIK) sebagai perluasan dari SID untuk memonitor penyaluran dana oleh lembaga keuangan. Namun, SLIK OJK menghadapi berbagai keluhan terkait akses yang lambat dan proses verifikasi yang lama. Memahami komentar-komentar pengguna penting untuk dilakukan karena berisi ulasan pengguna ketika menggunakan layanan, sehingga dapat dimanfaatkan oleh pengembang untuk meningkatkan kualitas pelayanan yang dimiliki. Akan tetapi, pengolahan informasi dari banyaknya komentar tidak memungkinkan untuk dilakukan secara manual. Oleh karena itu, penelitian ini menerapkan analisis sentimen menggunakan algoritma Naïve Bayes yang dioptimasi dengan Information Gain dan SMOTE. Data yang digunakan adalah 233 komentar pengguna layanan SLIK OJK dari Januari hingga Maret 2024. Model terbaik, Opt2, menunjukkan akurasi 77.65% dan F1-score 70.27%. Metode optimasi meningkatkan nilai precision, recall, dan F1-score masing-masing sebesar 29.45%, 13.78%, dan 16.71%, namun menurunkan akurasi sebesar 0.24%.

Kata kunci : SLIK Ojk, Analisis Sentimen , *Information Gain*, *Naïve Bayes*, SMOTE

1. PENDAHULUAN

Transformasi digital telah merangsang inovasi di sektor keuangan, khususnya melalui financial technology (fintech) seperti pinjaman online (pinjol), yang mengubah akses tradisional ke layanan keuangan[1]. Bank Indonesia Checking (BI Checking) atau Sistem Informasi Debitur (SID) menjadi semakin penting sebagai mekanisme verifikasi data nasabah dan mitigasi risiko bagi penyedia layanan fintech pinjol. SID Bank Indonesia (BI) mengelola data kredit yang sering dipakai oleh lembaga keuangan. Pada 27 April 2017, Otoritas Jasa Keuangan (OJK) memperkenalkan Sistem Layanan Informasi Keuangan (SLIK), perluasan dari SID, untuk memonitor penyaluran dana oleh lembaga keuangan kepada masyarakat. OJK memiliki tanggung jawab dan wewenang untuk mengawasi sektor perbankan dan lembaga jasa keuangan, menggantikan peran Bank Indonesia.

Sistem Layanan Informasi Keuangan (SLIK) memberikan skor pada setiap nasabah debitur berdasarkan catatan kreditnya. Skor kredit yang diberikan dihitung dari 1-5. Dalam sistem BI Checking oleh OJK, skor kredit dinilai dari 1 hingga 5. Berikut ini adalah penjelasan kategori kredit berdasarkan skornya:

- a. Skor 1: Kredit Lancar, artinya seperti seseorang yang selalu tepat waktu dalam memenuhi kewajibannya, membayar cicilan setiap bulan beserta bunganya hingga lunas tanpa pernah terlambat.
- b. Skor 2: Kredit DPK atau Kredit dalam Perhatian Khusus, artinya seperti seseorang yang terlambat membayar cicilan kredit antara 1 hingga 90 hari.

- c. Skor 3: Kredit Tidak Lancar, artinya seperti seseorang yang terlambat membayar cicilan kredit antara 91 hingga 120 hari.
- d. Skor 4: Kredit Diragukan, artinya seperti seseorang yang terlambat membayar cicilan kredit antara 121 hingga 180 hari.
- e. Skor 5: Kredit Macet, artinya seperti seseorang yang terlambat membayar cicilan kredit lebih dari 180 hari.

Bank akan menolak pengajuan kredit dari calon debitur yang memiliki skor 3, 4, atau 5 dalam BI Checking, yang berarti mereka termasuk dalam daftar hitam BI Checking. Sebab bank sama sekali tak mau ambil risiko kalau nantinya kredit yang diberikan bermasalah atau Non Performing Loan (NPL). Berdasarkan databoks.katadata.co.id pada bulan Juni 2023 tercatat nilai kredit macet pinjaman online dari nasabah debitur mencapai kurang lebih 1 triliun 345 miliar yang dipisahkan berdasarkan kelompok usia <19 tahun sampai usia >54 tahun. Hal ini akan berdampak pada skor BI Checking yang dimiliki pada nasabah debitur yang mengalami kemacetan dan proses pinjaman yang tidak akan disetujui dikarenakan skor BI Checking yang tidak mumpuni. Selain itu, menyatakan kemanaan dan kerahasiaan SLIK menimbulkan keraguan di kalangan masyarakat dan bank daerah. Hal ini disebabkan oleh laporan dari situs berita online Kompas.com yang mengungkapkan praktik jual beli data nasabah bank yang seharusnya bersifat rahasia, ada beberapa kasus juga seperti 5 orang pendaftar kerja tidak bisa mengikuti tes kerja di perusahaan karena Skor Kredit SLIK Kol 5, untuk pelayanan SLIK sendiri sangat susah karena website sering down dan permasalahan

susahnya mendapat antrian online dari SLIK itu sendiri. Permasalahan-permasalahan tersebut menimbulkan komentar-komentar yang diungkapkan oleh masyarakat pada media sosial seperti X.

Sebuah penelitian tentang analisis sentimen sebelumnya telah dilakukan oleh [2], di mana mereka menganalisis sentimen masyarakat mengenai pengembangan Pulau Rinca dengan menggunakan algoritma *Support Vector Machine* (SVM) dan *Logistic Regression* (LR). Model yang dikembangkan oleh peneliti ini mencapai akurasi sebesar 87% dan *F1-score* sebesar 81%. Namun, penelitian ini memiliki kelemahan karena penggunaan data yang tidak seimbang dalam dataset, yang dapat mempengaruhi hasil yang diperoleh. Selanjutnya, [3]. Melakukan perbandingan model untuk kampanye “anti-LGBT” di Indonesia dengan menggunakan algoritma *Naïve Bayes* (NB), *Decision Tree* (DT), dan *Random Forest* (RF). Hasilnya menunjukkan bahwa algoritma NB memiliki kinerja paling optimal dengan akurasi sebesar 83,43%. Namun, kelemahan dari penelitian ini adalah tidak mengubah kata yang tidak baku menjadi baku, yang dapat mempengaruhi nilai akurasi yang diperoleh.

Berdasarkan pembahasan tersebut, penulis mengusulkan pembangunan model klasifikasi teks untuk mengidentifikasi layanan SLIK OJK yang perlu ditingkatkan atau dipertahankan menggunakan algoritma *Naïve Bayes*. Penelitian ini akan melakukan normalisasi kata untuk mengubah kata tidak baku menjadi baku, dan menerapkan optimasi *Information Gain* untuk meningkatkan akurasi dan *F1-score*. Pendekatan ini sejalan dengan penelitian yang menerapkan metode optimasi seleksi fitur *Information Gain* dalam analisis sentimen maskapai penerbangan, yang menunjukkan peningkatan akurasi sebesar 6,6% dan *F1-score* sebesar 12,9% [4]. Selain itu, metode SMOTE akan digunakan untuk mengatasi ketidakseimbangan dataset, yang diharapkan dapat meningkatkan akurasi seperti yang tercatat dalam penelitian sebelumnya dengan peningkatan sebesar 1,918%. Penelitian ini diharapkan dapat membantu masyarakat memahami SLIK OJK berdasarkan ulasan pengguna di X.

2. TINJAUAN PUSTAKA

2.1. SLIK Ojk (Sistem Informasi Layanan Keuangan)

Sistem Layanan Informasi Keuangan (SLIK) yang dikelola oleh OJK mendukung pengawasan dan layanan informasi keuangan. SLIK memperlancar penyediaan dana, manajemen risiko kredit, penilaian kualitas debitur, serta verifikasi kerja sama dengan pihak ketiga. SLIK memperluas cakupan iDeb dengan mencakup lembaga keuangan bank, pembiayaan, dan non-bank. Setiap lembaga ini wajib melaporkan data debitur mereka ke Sistem Informasi Debitur (SID). Selain itu, SLIK juga melaporkan informasi tentang fasilitas penyediaan dana, data angsuran, dan data

terkait lainnya dari berbagai lembaga keuangan dan pihak-pihak yang relevan [5].

2.2. Analisis Sentimen

Analisis sentimen adalah proses di mana kita mengambil informasi dari suatu entitas dan secara otomatis menentukan apakah informasi dalam teks tersebut bersifat positif, negatif, atau netral. Proses ini melibatkan penggunaan metode pemrosesan teks yang memungkinkan kita untuk menganalisis dan memahami teks secara otomatis. [6]

2.3. Naïve Bayes

Naïve Bayes adalah metode klasifikasi yang didasarkan pada *Teorema Bayes*, dengan asumsi bahwa setiap pasangan fitur dalam kelas bersifat independen satu sama lain. Algoritma *Naïve Bayes* adalah metode klasifikasi yang efisien dan sederhana untuk menganalisis sentimen [7]. Rumus atau persamaan dari *Teorema Bayes* dapat dihitung menggunakan Persamaan.

$$P(M|N) = \frac{P(N|M)P(M)}{P(N)} \quad (1)$$

Keterangan

- N : Data sampel dengan label atau kelas yang belum teridentifikasi
- M : Hipotesis bahwa data N merupakan sebuah data yang spesifik dari suatu kelas
- $P(M|N)$: *Posterior probability* atau peluang terjadinya hipotesis M, setelah diberikan kondisi N
- $P(M)$: *Prior probability* atau peluang hipotesis M terjadi
- $P(N|M)$: Peluang terjadinya N, setelah diberikan kondisi hipotesis M
- $P(N)$: *Predictor prior probability* atau peluang terjadinya N

2.4. Information Gain

Information Gain adalah teknik seleksi fitur yang mengukur sejauh mana kehadiran atau ketidakhadiran suatu kata mempengaruhi pembuatan keputusan klasifikasi yang akurat dalam suatu kelas data. Tujuannya adalah untuk mengurangi ukuran dimensi dan meningkatkan akurasi klasifikasi. Rumus-rumus yang digunakan untuk menghitung nilai *Information Gain* [8] dapat ditemukan pada Persamaan 2, 3, dan 4.

$$\text{Info}(D) = - \sum_{i=1}^c p_i \log_2(p_i) \quad (2)$$

Persamaan 2 mengacu pada tingkat pentingnya atribut terhadap sebuah kelas atau nilai *entropy* suatu kelompok data, yang dinotasikan sebagai D. Jumlah partisi atau bagian dari D direpresentasikan oleh c. p_i mewakili probabilitas *tuple* pada D yang dimasukkan ke dalam c_i .

$$\text{Info}_A(D) = - \sum_{j=1}^v \frac{|D_j|}{|D|} * \text{Info}(D_j) \quad (3)$$

$\text{Info}_A(D)$ pada Persamaan 3 mempresentasikan pengukuran berapa banyak informasi yang diberikan atribut A dalam memprediksi kelas atau label target pada kelompok data D. Total seluruh bagian atribut A direpresentasikan oleh v . $|D_j|$ mewakili total nilai data dalam bagian j, $|D|$ merupakan total nilai data.

$$\text{InfoGain}(A) = -\text{Info}(D) - |\text{Info}(A)| \quad (4)$$

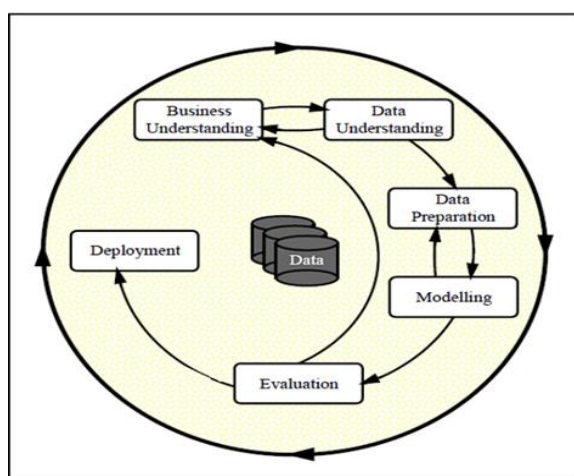
$\text{InfoGain}(A)$ pada Persamaan 4 merepresentasikan nilai dari *information gain* dari atribut A, sementara itu $\text{Info}(D)$ merepresentasikan jumlah keseluruhan nilai *entropy* pada kelompok data D, dan $\text{Info}(A)$ merepresentasikan nilai *entropy* dari atribut A.

2.5. SMOTE (Synthetic Minority Over-Sampling Technique)

SMOTE adalah teknik yang dapat mengatasi ketidakseimbangan distribusi data pada kelas minoritas dengan melakukan seleksi terhadap jumlah sampel hingga mencapai proporsi yang sama dengan kelas mayoritas. Menurut [9] SMOTE digunakan untuk mengatasi masalah overfitting pada imbalanced dataset. SMOTE bekerja dengan mencari tetangga terdekat dari k dalam setiap data yang termasuk dalam kelas minoritas, kemudian membuat data sintetik atau data buatan diantara data minoritas dengan data tetangga terdekat yang dipilih secara acak [8].

3. METODE PENELITIAN

Metodologi penelitian yang akan digunakan adalah CRISP-DM (*Cross Industry Standard Process for Data Mining*). Metodologi CRISP-DM memiliki beberapa tahapan dalam penerapannya, yaitu *business understanding*, *data understanding*, *data preparation*, *modelling*, *evaluation*, dan *deployment* [10].



Gambar 1 Tahapan CRIPS DM

Metode CRIPS-DM Sering digunakan pada Text Mining, Metode ini memiliki 6 tahapan, Yakni:

a. Bussines Understanding

Tahapan ini dimulai dengan menganalisis ulasan pelayanan SLIK OJK di X menggunakan algoritma *Naive Bayes* dan *Information Gain*. Tujuannya adalah menghasilkan model yang dapat mengidentifikasi ulasan pengguna SLIK OJK, dengan menggunakan tools seperti Python, *Tweet Harvest*, NLTK, Pandas, *Scikit-learn*, dan Sastrawi.

b. Data Understanding

Pada tahap ini, data ulasan pengguna SLIK OJK di X dikumpulkan dan dieksplorasi untuk menentukan kategori label positif dan negatif. Selanjutnya, data diverifikasi untuk memastikan kualitas dan kesesuaiannya dengan objek penelitian.

c. Data Preparation

Tahap *data preparation* meliputi penghapusan duplikasi dan ulasan netral, pelabelan manual, dan pembersihan data. Kemudian, dilakukan *case folding* untuk menyamakan format, *tokenization* untuk memisahkan kata, *normalization* untuk mengubah kata tidak baku, serta *stopword removal* dan *stemming* untuk menghilangkan kata tidak signifikan dan imbuhan.

d. Modeling

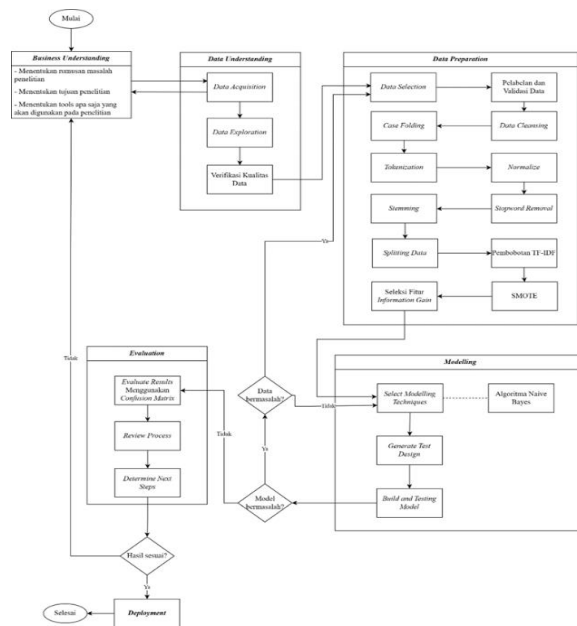
Tahap ini dimulai dengan memilih algoritma *Naive Bayes* sebagai teknik pemodelan, diikuti dengan menentukan test design berupa *confusion matrix*. Selanjutnya, model dibangun dan diuji. Hasil dari tahap ini akan dievaluasi pada tahap berikutnya.

e. Evaluation

Tahap ini melibatkan evaluasi model berdasarkan *confusion matrix*, diikuti dengan meninjau seluruh proses untuk mengidentifikasi keberhasilan dan kekurangan. Selanjutnya, analisis hasil evaluasi dan review digunakan untuk menentukan langkah selanjutnya untuk model.

f. Deployment

Pada tahap ini dilakukan pendokumentasian penelitian dalam bentuk jurnal publikasi dan skripsi. Setiap bagian dari hasil penelitian dibahas dan disajikan pada tahapan ini. Adapun hasil analisis sentimen divisualisasikan dalam bentuk *word cloud* sehingga dapat diketahui informasi apa saja yang terkandung dalam ulasan-ulasan yang diberikan oleh pengguna SLIK OJK di X dalam rentang waktu pada bulan Januari 2024 sampai dengan bulan Maret 2024.



Gambar 2 Kerangka Penelitian

4. HASIL DAN PEMBAHASAN

4.1. Business Understanding

Tahapan *business understanding* dilakukan pemahaman masalah yang berdasarkan pada latar belakang dan tujuan penelitian. Permasalahan pada SLIK ojk adalah pelayanan, kelebihan, dan kekurangan yang dimiliki layanan menyebabkan banyaknya komentar yang timbul dari pengguna. Komentar-komentar pengguna bermanfaat bagi pihak pengembang layanan dalam mengembangkan layanan yang dimiliki sesuai dengan kebutuhan penggunanya. Akan tetapi, begitu banyaknya komentar mengakibatkan proses pengolahan informasi tidak memungkinkan jika dilakukan secara manual karena akan memakan waktu yang lama. Sehingga diperlukan metode atau pendekatan yang dapat memudahkan proses pengolahan informasi. Oleh karena itu, tujuan penelitian ini adalah menerapkan analisis sentimen menggunakan pendekatan *machine learnin*, serta mengetahui kinerja dari pendekatan tersebut dalam melakukan klasifikasi komentar pengguna layanan SLIK ojk. Pendekatan *machine learning* yang digunakan adalah *Naïve Bayes*, yang dioptimalkan dengan *Information Gain* dan *SMOTE* untuk meningkatkan akurasi serta *F1-Score* yang diperoleh. Penerapan *SMOTE* juga bertujuan untuk mengatasi *imbalanced dataset* yang digunakan pada penelitian. Model dibangun menggunakan *tools* Google Colaboratory dengan Bahasa Python dan *library* Sastrawi, NLTK, *Scikit-learn*, serta *Pandas*.

4.2. Data Understanding

Tahap *data understanding* dimulai dengan pengumpulan data awal (*data acquisition*) yang akan digunakan pada penelitian, kemudian dilakukan eksplorasi data untuk mendapatkan informasi dari data dan menentukan kategori data yang akan digunakan.

Adapun tahapan dari *data understanding* ini adalah sebagai berikut:

4.2.1. Data Acquisition

Proses *data acquisition* atau pengumpulan data komentar pengguna Layanan SLIK ojk dilakukan dengan bantuan *library* *Tweet Harvest*. Data yang digunakan yaitu ulasan pengguna SLIK ojk berbahasa Indonesia pada periode Januari 2024 sampai dengan Maret 2024. Jumlah data komentar pengguna Layanan SLIK ojk yang dikumpulkan dari proses ini dalam rentang waktu Januari 2024 sampai dengan Maret 2024 adalah sebanyak 316 data. Data tersebut belum terstruktur dan mengandung banyak gangguan seperti duplikat, emotikon, tanda baca, angka, dan kata-kata yang tidak relevan dalam proses klasifikasi.

4.2.2. Data Exploration

Tahapan *data exploration* dilakukan untuk mengetahui informasi yang ada pada data, kemudian ditentukan kategori apa yang digunakan sesuai dengan tujuan penelitian. Tahapan *data exploration* ini didapatkan informasi sebagai berikut:

- Semua data yang dikumpulkan dari proses *data acquisition* berupa ulasan mengenai SLIK ojk.
- Jumlah data komentar yang didapatkan sebanyak 316 data komentar.
- Semua komentar berbahasa Indonesia. Kategori yang akan digunakan pada data merupakan kategori positif dan kategori negatif. Pengkategorian akan dilakukan secara manual.

4.2.3. Verifikasi Kualitas Data

Tahapan verifikasi ini dilakukan secara manual dengan melihat setiap ulasan pengguna dan memastikan ulasan tersebut sesuai dengan yang dibutuhkan. Hasil dari tahap ini menunjukkan bahwa data yang dikumpulkan telah memenuhi kebutuhan, yaitu ulasan pengguna tentang SLIK OJK. sehingga dapat dilanjutkan ke tahapan berikutnya untuk menyiapkan data sebelum proses klasifikasi.

4.3. Data Preparation

Tahapan persiapan data dilakukan untuk mengubah data yang tidak terstruktur menjadi data terstruktur, sehingga mempermudah proses klasifikasi. Tahapan *data preparation* ini terdiri dari 12 proses yaitu *data selection*, pelabelan dan validasi data, *data cleansing*, *case folding*, *tokenization*, *normalize*, *stopword removal*, *stemming*, *splitting data*, pembobotan *TF-IDF*, *SMOTE*, dan seleksi fitur *Information Gain*. Adapun tahapan ini dilaksanakan sebagai berikut:

4.3.1. Data Selection

Tahapan ini dilakukan dengan cara menyeleksi data ulasan pengguna dengan ulasan yang sama atau duplikasi data. Data awal yang diperoleh sebanyak 316 data ulasan, terdapat 83 data ulasan yang sama atau duplikasi data d, sehingga total data setelah dilakukan

proses seleksi ini sebanyak 233 data komentar pengguna Layanan SLIK ojk.

4.3.2. Pelabelan dan validasi Data

Proses klasifikasi memerlukan sebuah kategori atau label yang dapat membantu model yang dikembangkan dalam mengenali ulasan-ulasan pengguna. Kategori yang digunakan merupakan positif dan negatif, dilakukan secara manual dengan mengamati setiap data komentar dari hasil tahapan sebelumnya. Adapun jumlah data ulasan dengan label positif 168 data dan negatif 65 data.

4.3.3. Data Cleansing

Proses *Data Cleansing* sampai Proses *Stemming* akan di sajikan dalam Tabel 1.

Tabel 1. Tahap Data Cleansing sampai Stemming

Tahap	Hasil
Komentar Awal	SLIK OJK emang selelet itu ngab buat diakses?
Data Cleansing	SLIK OJK emang selelet itu ngab buat diakses
Case Folding	slik ojk emang selelet itu ngab buat diakses
Tokenization	['slik', 'ojk', 'emang', 'selelet', 'itu', 'ngab', 'buat', 'diakses']
Normalize	['slik', 'ojk', 'emang', 'selelet', 'itu', 'ngab', 'buat', 'diakses']
Stopword Removal	['slik', 'ojk', 'emang', 'selelet', 'buat', 'diakses']
Stemming	['slik', 'ojk', 'emang', 'lelet', 'buat', 'akses']

4.3.4. Splitting Data

Tahapan ini dilakukan dengan cara memisahkan data menjadi data latih (*data training*) dan data uji (*data testing*). Data dibagi berdasarkan skenario yang telah dibuat sebelumnya pada Tabel 2.

Tabel 2. Splitting Data

Skenario	Data Latih (<i>Data Training</i>)	Data Uji (<i>Data Testing</i>)
50:50	116	117
60:40	139	94
70:30	163	70
80:20	186	47
90:10	209	24

4.3.5. TF-IDF dan Information Gain

Setelah melakukan pemisahan data (*splitting data*), tahapan berikutnya adalah pembobotan TF-IDF. Pada tahap ini, bobot diberikan pada setiap kata berdasarkan kemunculannya dalam dokumen dengan menggunakan *library scikit-learn*. Selanjutnya, dilakukan seleksi fitur menggunakan *Information Gain* untuk memilih fitur yang paling berpengaruh. Fitur dipilih berdasarkan skor yang diperoleh dan dibandingkan dengan nilai ambang batas yang telah ditentukan, yaitu *Threshold* 0.003. Hasil dari seleksi fitur ini dapat dilihat pada Tabel 3.

Tabel 3. Tahap TF-IDF Dan *Information Gain*

Skenario	TF-IDF	<i>Information gain</i>
50:50	942	942
60:40	1047	1047
70:30	1175	639
80:20	1264	696
90:10	1320	708

4.3.6. SMOTE

Tahapan selanjutnya yang dilakukan setelah memberikan bobot dengan TF-IDF adalah penerapan SMOTE. Tahapan ini dilakukan dengan cara menyeimbangkan *dataset* yang digunakan pada penelitian yaitu data dengan label positif dan negatif. Adapun hasil dari penerapan SMOTE ini dapat dilihat pada Tabel 4.

Tabel 4. Tahap SMOTE

Skenario	Sebelum Penerapan SMOTE		Setelah Penerapan SMOTE	
	Positif	Negatif	Positif	Negatif
50:50	77	39	77	77
60:40	97	42	97	97
70:30	116	47	116	116
80:20	136	50	136	136
90:10	154	55	154	154

4.4. Modeling

Tahap ini bertujuan untuk melakukan pemodelan menggunakan data yang sudah disiapkan pada tahap *data preparation*. Tahap pertama yang dilakukan adalah *select modeling techniques* atau pemilihan teknik pemodelan. Penelitian ini menggunakan model algoritma *Naïve Bayes*. Kemudian dilanjutkan dengan menentukan desain pengujian model. Setelah itu melakukan pembangunan model untuk pelatihan. Selanjutnya melakukan *testing* model untuk menilai kinerja dari model yang sudah dihasilkan. Adapun tahapan *modeling* yang dilakukan pada penelitian ini adalah sebagai berikut:

a. Select Modeling Techniques

Proses ini dilakukan dengan cara memilih teknik pemodelan yang akan digunakan yaitu menggunakan model algoritma *Naïve Bayes*. Model algoritma *Naïve Bayes* di-*import* dari *library scikit-learn*. Setelah itu di dapatkan 10 Model Skenario 5 Model Ori 5 Model Opt.

- Model Ori1, dilatih menggunakan skenario *splitting data* 50:50 sebelum penerapan SMOTE dan sebelum seleksi fitur *Information Gain*.
- Model Ori2, dilatih menggunakan skenario *splitting data* 60:40 sebelum penerapan SMOTE dan sebelum seleksi fitur *Information Gain*.
- Model Ori3, dilatih menggunakan skenario *splitting data* 70:30 sebelum penerapan SMOTE dan sebelum seleksi fitur *Information Gain*.

- Model Ori4, dilatih menggunakan skenario *splitting data* 80:20 sebelum penerapan SMOTE dan sebelum seleksi fitur *Information Gain*.
 - Model Ori5, dilatih menggunakan skenario *splitting data* 90:10 sebelum penerapan SMOTE dan sebelum seleksi fitur *Information Gain*.
 - Model Opt1, dilatih menggunakan skenario *splitting data* 50:50 setelah penerapan SMOTE dan setelah seleksi fitur *Information Gain*.
 - Model Opt2, dilatih menggunakan skenario *splitting data* 60:40 setelah penerapan SMOTE dan setelah seleksi fitur *Information Gain*.
 - Model Opt3, dilatih menggunakan skenario *splitting data* 70:30 setelah penerapan SMOTE dan setelah seleksi fitur *Information Gain*.
 - Model Opt4, dilatih menggunakan skenario *splitting data* 80:20 setelah penerapan SMOTE dan setelah seleksi fitur *Information Gain*.
 - Model Opt5, dilatih menggunakan skenario *splitting data* 90:10 setelah penerapan SMOTE dan setelah seleksi fitur *Information Gain*.
- b. *Generate Test Design*
- Pemilihan metode pengujian dilakukan dalam proses ini untuk menguji performa dari model yang telah dibuat. Penelitian ini menggunakan metode *confusion matrix* untuk mengevaluasi kinerja model berdasarkan nilai *accuracy*, *precision*, *recall*, dan *F1-score* yang dihasilkan.

4.5. Evaluation

Tahap ini dilakukan dengan cara mengevaluasi secara menyeluruh dari setiap proses atau langkah yang dijalankan pada tahapan sebelumnya. Proses dalam tahapan ini dimulai dari evaluate result, review process, dan determine next steps. Adapun hasil prosesnya bisa dilihat di Tabel 5.

Tabel 5. Pengaruh Optimasi Pada Nilai *Confusion Matrix*

Model Ori (Sebelum Optimasi)			Model Opt (Setelah Optimasi)			Dampak
Model	Pengukuran	Hasil	Model	Pengukuran	Hasil	
Ori1	Akurasi	77.78%	Opt1	Akurasi	71.79%	-5.99%
	Recall	50.00%		Recall	59.89%	+9.89%
	Precision	38.89%		Precision	59.63%	+20.74%
	F1-Score	43.75%		F1-Score	59.75%	+16.00%
Ori2	Akurasi	75.53%	Opt2	Akurasi	77.66%	+2.13%
	Recall	50.00%		Recall	70.51%	+20.51%
	Precision	37.77%		Precision	69.94%	+32.17%
	F1-Score	43.03%		F1-Score	70.21%	+27.18%
Ori3	Akurasi	74.29%	Opt3	Akurasi	72.86%	-1.43%
	Recall	50.00%		Recall	54.49%	+4.49%
	Precision	37.14%		Precision	59.52%	+22.38%
	F1-Score	42.62%		F1-Score	53.74%	+11.12%
Ori4	Akurasi	68.09%	Opt4	Akurasi	63.83%	-4.26%
	Recall	05.00%		Recall	52.19%	+2.198%
	Precision	34.04%		Precision	53.37%	+19.33%
	F1-Score	40.51%		F1-Score	51.07%	+10.56%
Ori5	Akurasi	58.33%	Opt5	Akurasi	66.67%	+8.34%
	Recall	50.00%		Recall	60.00%	+10.00%
	Precision	29.17%		Precision	81.82%	+52.65%
	F1-Score	36.84%		F1-Score	55.56%	+18.72%

Berdasarkan Tabel 5 pengaruh peningkatan paling signifikan optimasi Nilai pengukuran *Confusion Matrix* ada di Skenario Model Ori2 Ke Skenario Model Opt2 untuk peningkatan akurasi sebesar 2.13%, *recall* 20.51%, *precision* 32.17% dan *F1-Score* 27.18%.

Setelah itu, dilakukan proses peninjauan untuk memastikan bahwa langkah-langkah yang telah diambil sesuai dengan tujuan penelitian. Kemudian, tahap selanjutnya ditetapkan dalam proses menentukan langkah berikutnya. Penelitian ini telah menyelesaikan semua tahapan dan menghasilkan model yang mampu mengklasifikasikan komentar pengguna, sehingga dapat dilanjutkan ke tahap penerapan (*deployment*).

4.6. Deployment

Proses pemodelan telah berhasil dilakukan dan telah dievaluasi. Maka dari itu, Selanjutnya, dilakukan tahap dokumentasi dengan menyajikan setiap langkah yang telah dilakukan dalam bentuk laporan penelitian berupa skripsi dan publikasi jurnal. Adapun informasi dari hasil analisis yang didapatkan direpresentasikan dalam bentuk *word cloud* dan dapat dilihat pada Gambar 3 untuk komentar dengan label positif dan pada Gambar 4 untuk komentar dengan label negatif.



Gambar 3. *Wordcloud* Ulasan Positif

Berdasarkan Gambar 3 kata yang divisualisasikan dari komentar yang memiliki sentimen positif adalah kata “informasi” yang muncul menunjukkan bahwa layanan ini memberikan informasi kepada penggunanya dalam mengakses informasi keuangan, hal ini dikuatkan dengan adanya kata “informasi” dalam *word cloud*.



Gambar 4. *Wordcloud* Ulasan Negatif

Berdasarkan Gambar 4 kata yang divisualisasikan dari komentar negatif pengguna adalah kata “gagal” dan “jelek” menunjukkan bahwa layanan ini memberikan pelayanan yang kurang baik, serta hal ini juga berlaku pada kata “lama” yang muncul pada *word cloud*. menunjukkan bahwa layanan mengalami gangguan ketika digunakan oleh pengguna.

5. KESIMPULAN DAN SARAN

Penelitian ini mengoptimasi analisis sentimen pengguna layanan SLIK OJK menggunakan SMOTE dan seleksi fitur Information Gain dengan algoritma Naïve Bayes, menghasilkan model terbaik (Opt2) dengan akurasi 77.65% dan F1-score 70.27%. Optimasi meningkatkan precision sebesar 16.71%, recall 29.45%, dan F1-score 13.78%, meskipun akurasi rata-rata menurun sebesar 0.24%. Disarankan untuk menggunakan metode klasifikasi lain seperti SVM, KNN, atau DT, mencoba metode optimasi lainnya, menggunakan lebih banyak data dengan rentang waktu yang lebih panjang, serta mempertimbangkan data multibahasa untuk penelitian selanjutnya.

DAFTAR PUSTAKA

- [1] F. Novika, N. Septivani, and I. M. Indra P, "Pinjaman Online Ilegal Menjadi Bencana Sosial Bagi Generasi Milenial," *Management Studies and Entrepreneurship Journal (MSEJ)*, vol. 3, no. 3, pp. 1174–1192, Aug. 2022, doi: 10.37385/msej.v3i3.857.
- [2] T. H. Jaya Hidayat, Y. Ruldeviyani, A. R. Aditama, G. R. Madya, A. W. Nugraha, and M. W. Adisaputra, "Sentiment analysis of twitter data related to Rinca Island development using Doc2Vec and SVM and logistic regression as classifier," *Procedia Comput Sci*, vol. 197, pp. 660–667, 2022, doi: 10.1016/j.procs.2021.12.187.
- [3] M. I. Fikri, T. S. Sabrila, and Y. Azhar, "Perbandingan Metode Naïve Bayes dan Support Vector Machine pada Analisis Sentimen Twitter," *SMATIKA JURNAL*, vol. 10, no. 02, pp. 71–76, Dec. 2020, doi: 10.32664/smatika.v10i02.455.
- [4] A. B. P. Negara, H. Muhandi, and I. M. Putri, "Analisis Sentimen Maskapai Penerbangan Menggunakan Metode Naive Bayes dan Seleksi Fitur Information Gain," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 7, no. 3, p. 599, May 2020, doi: 10.25126/jtiik.2020711947.
- [5] ojk.go.id, "Sistem Layanan Informasi Keuangan (SLIK)," ojk.go.id. Accessed: May 24, 2024. [Online]. Available: <https://ojk.go.id/id/kanal/perbankan/Pages/Sistem-Layanan-Informasi-Kuangan-SLIK.aspx#>
- [6] N. C. Dang, M. N. Moreno-García, and F. De la Prieta, "Sentiment Analysis Based on Deep Learning: A Comparative Study," *Electronics (Basel)*, vol. 9, no. 3, p. 483, Mar. 2020, doi: 10.3390/electronics9030483.
- [7] A. B. P. Negara, H. Muhandi, and I. M. Putri, "Analisis Sentimen Maskapai Penerbangan Menggunakan Metode Naive Bayes dan Seleksi Fitur Information Gain," vol. 7, no. 3, pp. 599–606, 2020, doi: 10.25126/jtiik.202071947.
- [8] H. Hidayatullah and Y. Umaidah, "PENERAPAN NAÏVE BAYES DENGAN

- OPTIMASI INFORMATION GAIN DAN SMOTE UNTUK ANALISIS SENTIMEN PENGGUNA APLIKASI CHATGPT,” 2023.
- [9] V. Rupapara, F. Rustam, H. F. Shahzad, A. Mehmood, I. Ashraf, and G. S. Choi, “Impact of SMOTE on Imbalanced Text Features for Toxic Comments Classification Using RVVC Model,” *IEEE Access*, vol. 9, pp. 78621–78634, 2021, doi: 10.1109/ACCESS.2021.3083638.
- [10] Y. Mahmud, N. S. Shaeali, and S. Mutalib, “Comparison of Machine Learning Algorithms for Sentiment Classification on Fake News Detection,” *International Journal of Advanced Computer Science and Applications*, vol. 12, no. 10, 2021, doi: 10.14569/IJACSA.2021.0121072.