

IMPLEMENTASI NAÏVE BAYES DAN SUPPORT VECTOR MACHINE UNTUK ANALISIS SENTIMEN ULASAN PADA *GOOGLE PLAY*

Yoke Lucia Renica Rehatalanit¹, Marcelino Paul Edwil Longdong²,
Achmad Ramadhany³

^{1,2,3}Universitas Dirgantara Marsekal Suryadarma

¹ylrrehatalanit@gmail.com, ²edwillongdong@gmail.com, ³aramadhany03@gmail.com

ABSTRACT

Reviews on Google Play describes user sentiment towards the application according to the ratings and comments are given. In practice, there is often a discrepancy between the rating and the comments given, resulting in a biased sentiment, so it is necessary to analyze the review to find out the sentiment contained therein. In collecting data from the Google Play site using the Web Scraping technique with the google-play-scraper package from Python. Reviews that are successfully scraped then go through the preprocessing stage so that the data set is more structured. In the next stage, the data set is labeled based on the rating, and given a weight using TF-IDF. After classifying using the Naïve Bayes and Support Vector Machine methods, then evaluating using the confusion matrix, and validating using K-Fold Cross Validation. Research results using the Naïve Bayes method and Support Vector Machine for sentiment analysis on the Google Play website, the Naïve Bayes method produces 87.82% accuracy, 58.90% precision, 60.08% recall, while the Support Vector Machine method produces 90% accuracy .01% , precision 61.89%, recall 60.18%.

Keywords: Sentiment Analysis, Naïve Bayes, Support Vector Machine, Google Play, Web Scraping, TF-IDF, Confusion Matrix.

ABSTRAK

Data Ulasan pada Google Play menggambarkan sentimen pengguna kepada aplikasi sesuai *rating* dan komentar yang diberikan. Dalam prakteknya sering kali terjadi ketidaksesuaian antara *rating* dan komentar yang diberikan sehingga terjadi bias sentimen, sehingga perlu dilakukan analisis terhadap ulasan tersebut untuk mengetahui sentimen yang terkandung didalamnya. Dalam pengambilan data dari situs *Google Play* menggunakan teknik *Web Scraping* dengan package *google-play-scraper* dari *Python*. Ulasan yang berhasil di-*scraping* kemudian melalui tahap *preprocessing* agar *data set* lebih terstruktur. Tahap selanjutnya *data set* diberikan label berdasarkan *rating*, serta diberikan bobot menggunakan TF-IDF. Setelah dilakukan klasifikasi menggunakan metode *Naïve Bayes* dan *Support Vector Machine*, kemudian dilakukan evaluasi menggunakan *Confusion Matrix*, dan divalidasi menggunakan *K-Fold Cross Validation*. Hasil Penelitian menggunakan metode *Naïve Bayes* dan *Support Vector Machine* untuk analisis sentimen pada situs *Google Play*, pada metode *Naïve Bayes* menghasilkan *accuracy* 87,82%, *precision* 58,90%, *recall* 60,08%, sementara pada metode *Support Vector Machine* menghasilkan *accuracy* 90,01%, *precision* 61,89%, *recall* 60,18%.

Kata Kunci : Analisis Sentimen, *Naïve Bayes*, *Support Vector Machine*, *Google Play*, *Web Scraping*, TF-IDF, *Confusion Matrix*.

PENDAHULUAN

Sentimen atau opini publik sangat berpengaruh terhadap citra suatu perusahaan maupun produk (Bahtera, Vidyarini, & ..., 2019). Pada *Google Play* pengguna dapat memberikan opini atau penilaiannya terhadap suatu aplikasi

melalui ulasan atau review. Ulasan atau review pada *Google Play*, umumnya berisi rating bintang 1 sampai 5 dan juga komentar yang didalamnya terdapat masukan, serta keluhan terhadap sebuah aplikasi, baik bersifat positif, netral,

maupun negatif. Rating dan komentar yang diberikan oleh pengguna secara tidak langsung dapat mempengaruhi calon pengguna baru untuk menggunakan aplikasi tersebut. Rating merupakan gambaran umum sentimen pengguna terhadap sebuah aplikasi, sehingga rating sangat berpengaruh terhadap citra sebuah aplikasi (Masturoh, 2021). Namun, ada kalanya pemberian rating tidak sesuai dengan komentar yang diberikan, sehingga belum dapat menggambarkan tanggapan dari pengguna secara utuh sehingga diperlukan analisis sentimen terhadap teks ulasan (Faadilah, 2020).

Analisis sentimen merupakan metode pengolahan bahasa alami atau *natural language processing* (NLP) untuk memilah emosi menjadi positif, negatif dan netral yang terdapat didalam tulisan tertentu (Awaludin & Ridyustia Raveena, 2021). Analisis sentimen bertujuan untuk mengetahui sentimen yang terdapat dalam opini seseorang dan mengklasifikasikan berdasarkan emosi yang terkandung didalamnya. Dari manfaat dan efek yang diberikan analisis sentimen, banyak penelitian dan pengembangan aplikasi yang mengangkat topik ini (Masturoh, 2021).

Untuk melakukan analisis terhadap sentimen pada ulasan pengguna bukanlah hal yang mudah jika dilakukan secara manual dengan jumlah *data set* yang besar. Karena itu diperlukan algoritma *Naïve Bayes* dan *Support Vector Machine* sebagai metode pengklasifikasian teks dengan pendekatan *supervised learning* untuk melihat dan mengklasifikasikan sentimen pada ulasan yang diberikan oleh pengguna (Anjasmos, Marisa, & Istiadi, 2020). Pada penelitian ini, metode klasifikasi *text* yang digunakan adalah *Naïve Bayes* dan *Support Vector Machine*. Selanjutnya, akan dijabarkan tahap-tahap yang akan dilewati untuk

melakukan implementasi dan perbandingan algoritma *Naïve Bayes* dan *Support Vector Machine*, serta bagaimana mengukur kualitas hasil analisis menggunakan beberapa parameter yaitu akurasi, presisi dan *recall*. Penelitian ini berfokus untuk mengimplementasikan dan membandingkan performansi kedua metode tersebut dalam klasifikasi teks, guna mendapatkan algoritma yang tingkat akurasinya lebih tinggi dari algoritma lainnya dalam analisis sentimen data ulasan pada situs *Google Play*, dengan *data set* yang digunakan.

Metode *Naïve Bayes* adalah algoritma klasifikasi data yang berdasar pada teorema bayes. Metode ini dapat digunakan untuk klasifikasi data kualitatif maupun kuantitatif (Awaludin, Yasin, & Risyda, 2024). Selain itu metode ini tidak memerlukan data training dalam jumlah besar serta dapat digunakan untuk klasifikasi masalah biner dan juga klasifikasi multi kelas (Litbang, 2021). Metode ini juga dapat melakukan perhitungan yang relatif cepat dan efisien, terbukti dari penelitian Ahmad (2020) pada evaluasi *confusion matrix* dengan pembobotan TF-IDF 1,43 detik dan TF 3 detik. Sedangkan pada *K-Fold Cross Validation* rata-rata waktu dengan pembobotan TF 3,5 detik sementara TF-IDF 0,8 detik.

Support Vector Machine merupakan metode dalam *machine learning* (*supervised learning*) yang dimanfaatkan untuk klasifikasi dan regresi (Awaludin, 2015). Metode ini umum digunakan karena, sangat efektif pada data-data dengan batas kelas yang jelas serta kondisi saat fitur yang ada berjumlah lebih besar dari jumlah titik data yang ada. Selain itu *Support Vector Machine* mampu mendeteksi asosiasi kompleks terhadap data meski sedikit transformasi yang dilakukan (Hussein, 2021).

Data ulasan yang ada pada *Google Play* akan terus bertambah seiring pertumbuhan pengguna, oleh karena itu pengumpulan data dalam penelitian ini menggunakan teknik *web scraping*. Teknik *web scraping* merupakan metode yang tepat untuk mengumpulkan data dalam jumlah besar secara otomatis sesuai kebutuhan (Flores et al., 2020).

Penelitian Kurniawan (2017) berjudul Implementasi *Text Mining* Pada Analisis Sentimen Pengguna Twitter Terhadap Media Mainstream Menggunakan *Naïve Bayes Classifier* Dan *Support Vector Machine*. Pada Metode *Naïve Bayes* didapatkan nilai akurasi 95,8% pada TV One, 97,8% pada Kompas TV dan 91% pada Metro TV. Sedangkan pada metode *Support Vector Machine* didapatkan nilai akurasi 97,9% pada TV One, 99,3% pada Kompas TV dan 99,1% pada Metro TV. Dari segi akurasi dapat disimpulkan bahwa performa metode *Support Vector Machine* lebih baik dalam mengklasifikasi data.

Penelitian lain dilakukan oleh Ilmawan dan Mude (2020) dengan judul Perbandingan Metode Klasifikasi *Support Vector Machine* dan *Naïve Bayes* untuk Analisis Sentimen pada Ulasan Tekstual di *Google Play Store*. Dari penelitian tersebut disimpulkan bahwa *Support Vector Machine* memiliki nilai akurasi lebih baik dibanding *Naïve Bayes* untuk mengklasifikasikan ulasan tekstual berbahasa Indonesia pada *Google Play Store*. Nilai akurasi yang diperoleh pada metode *Support Vector Machine* sebesar 81,46% dan *Naïve Bayes* sebesar 75,41%.

Ahmad (Ahmad, 2020) Dalam penelitiannya berjudul Studi Perbandingan Metode Analisis *Naïve Bayes Classifier* Dengan *Support Vector*

Machine Untuk Analisis Sentimen (Studi Kasus: Tweet Berbahasa Indonesia Tentang Covid-19), dengan pembobotan TF-IDF pada metode *Support Vector Machine* nilai akurasi sebesar 83%. Sedangkan pada metode *Naïve Bayes Classifier* memiliki nilai akurasi 82,3%. Pada evaluasi dengan *K-Fold Cross Validation* dan pembobotan TF-IDF, metode *Support Vector Machine* memiliki akurasi 82,6% dan metode *Naïve Bayes* memiliki akurasi 82%. Hasil penelitian diatas menunjukkan bahwa metode *Support Vector Machine* memiliki nilai akurasi yang lebih baik dari algoritma *Naïve Bayes Classifier*, pada pembobotan TF-IDF serta saat ditambahkan *K-Fold Cross Validation*. Karena itu penelitian ini menggunakan *Naïve Bayes* dan *Support Vector Machine* untuk mengklasifikasikan sentimen ulasan pada *Google Play*.

METODE PENELITIAN

Dalam sebuah penelitian diperlukan metode agar dapat terlaksana secara terstruktur. Metode yang digunakan penulis adalah metode eksperimen. Berikut merupakan tahap-tahap yang dilakukan antara lain:

1. Pengumpulan Data

Pada penelitian ini data bersumber dari *website Google Play* dengan alamat <https://play.google.com/store/apps/details?id=id.co.bri.brimo>. Data yang didapat berupa *review* atau ulasan pengguna aplikasi *mobile banking* yang dikumpulkan mulai tanggal 1 Januari 2022 – 27 Juni 2022. Data dikumpulkan menggunakan teknik *web scraping*. Data di-*scraping* menggunakan bahasa pemrograman *Python* dengan *package google-play-scraper* (Python Software Foundation, 2021). Berikutnya data disimpan menggunakan format (.csv).

2. Pengolahan Data Awal

Pada tahap ini akan dilakukan beberapa proses yaitu :

A. *Text Preprocessing*

Tahap ini merupakan tahap untuk mengubah data yang belum terstruktur menjadi data terstruktur. Tahap *text preprocessing* memiliki beberapa proses yang umum dilakukan, yaitu *case folding*, *cleansing*, *tokenize*, *normalization*, *stopword removal*, dan *stemming*.

B. Pelabelan

Pada tahap ini dilakukan pemberian label kepada *data set* berdasarkan *rating* pada ulasan. Label terbagi menjadi tiga kelas, yaitu positif, netral, dan negatif. Ulasan yang memiliki *rating* 4 dan 5 diberikan label positif, ulasan yang memiliki *rating* 3 diberikan label netral, sedangkan ulasan yang memiliki *rating* 1 dan 2 diberikan label negatif. *Data set* yang sudah diberikan label, kemudian disajikan dalam bentuk grafik serta visualisasi *wordcloud* berdasarkan masing-masing kelas.

C. Pembobotan

Pada tahap ini *data set* yang masih berbentuk teks ditransformasikan kedalam bentuk numerik atau angka, agar dapat dibaca oleh program komputer. Tahap ini menggunakan metode pembobotan *term frequency - inverse document frequency* (TF-IDF), dimana nilai bobot akan sesuai dengan jumlah kata dalam teks.

3. Metode Yang Diusulkan

Berdasarkan penelitian terdahulu yang sejenis maka, penulis mengusulkan metode *Naïve Bayes* dan *Support Vector Machine* untuk melakukan analisis sentimen pada penelitian ini.

4. Eksperimen dan Pengujian Metode

Pada tahap ini *data set* yang sudah diberikan bobot, kemudian akan dibagi kedalam data *training* dan data *testing*. Dalam penelitian ini, data *training* dan data *testing* dibagi kedalam tiga rasio perbandingan, yaitu 90% : 10%, 80% : 20%, dan 70% : 30%. Hal ini dilakukan agar mendapatkan rasio dengan tingkat akurasi terbaik. Setelah data *training* dan data *testing* dibagi kemudian dilakukan implementasi algoritma *Naïve Bayes* dan *Support Vector Machine* untuk klasifikasi data uji berdasarkan data latih.

5. Evaluasi dan Validasi Hasil

Pada tahap ini, hasil pengujian dengan nilai akurasi terbaik dari rasio data *training* dan data *testing* masing-masing metode akan dievaluasi menggunakan *confusion matrix*, untuk dan divalidasi menggunakan *K-Fold Cross Validation*, untuk mengetahui nilai maksimal *accuracy*, *recall*, dan *precision* dari data yang telah diuji.

PEMBAHASAN

1. Pengumpulan Data

Data ulasan yang digunakan pada penelitian ini bersumber dari situs *Google Play* dengan alamat <https://play.google.com/store/apps/details?id=id.co.bri.brimo>. Data ulasan di-*scraping* menggunakan *software google collab*. *Scraping* data ulasan ini juga memanfaatkan *library pandas*, *numpy*, dan *package google-play-scraper* dari bahasa pemrograman *Python*. Code untuk proses *scraping* data pada Gambar 1 dibawah ini.

Gambar 1 Implementasi *Scraping* Data Ulasan

[illegible]

Gambar 2 Hasil *Scraping* Data Ulasan

Tabel 1 Deskripsi Variabel *Scraping*
Data Ulasan

Variabel	Deskripsi
content	Pendapat atau opini pengguna
score	Nilai atau <i>rating</i> pengguna
at	Tanggal dan jam ulasan diberikan

Dalam *Text Preprocessing* terdapat beberapa langkah yang umum dilakukan yaitu *case folding*, *tokenizing*, *stopword removal*, *normalization* dan *stemming*.

Tahap yang pertama dijalankan dalam *text preprocessing* adalah *case folding*. *Case Folding* bertujuan untuk membuat seluruh huruf yang terdapat pada *data set* bagian ulasan menjadi *lower text* atau huruf kecil. *Code* dari tahap ini dapat dilihat pada Tabel 2 merupakan contoh hasil penerapan dari tahap *case folding*.

Sebelum	Sesudah
Mantap segala sesuatu jadi gampang ,,	mantap segala sesuatu jadi gampang ,,

Tahap berikutnya ialah *Cleansing*, pada tahap ini dilakukan pembersihan pada *data set* terhadap elemen-elemen yang tidak dibutuhkan, seperti *ascii*, angka, *URL*, *hastag*, *tab*, *new line*, dan *back slice*. Hasil dari tahap *cleansing* dapat dilihat pada Tabel 3 dibawah ini.

Sebelum	Sesudah
Mantap segala sesuatu jadi gampang ,,	mantap segala sesuatu jadi gampang

Tahap *tokenizing* diaplikasikan untuk memecah menjadi token-token atau penggalan kata. Tahap *tokenizing* ini memanfaatkan *package NLTK tokenize* dengan *library word_tokenize*. Hasil implementasi dari tahap *tokenizing* dapat dilihat pada Tabel 4 dibawah ini.

Sebelum	Sesudah
---------	---------

mantap segala sesuatu jadi gampang	['mantap', 'segala', 'sesuatu', 'jadi', 'gampang']
--	--

4). Normalization

Normalization dilakukan untuk memperbaiki kata yang memiliki penulisan yang salah (*typo*) atau kata yang penulisannya disingkat. *Normalization* berguna untuk menyeragamkan kata yang mempunyai makna yang sama namun berbeda penulisannya. Pada penelitian ini digunakan kamus normalisasi yang bersumber dari Kamus NLP (*Natural Language Processing*) Bahasa Indonesia Resource di Github (Owen, 2020). Tabel 5 merupakan hasil dari tahap normalisasi.

Tabel 5 Hasil *Normalization*

Sebelum	Sesudah
['knp', 'tidak', 'bisa', 'daftar']	['kenapa', 'tidak', 'bisa', 'daftar']

5). Stopword Removal

Tahap *stopword removal* bertujuan untuk menghapus kata-kata umum yang banyak digunakan tapi tidak memberikan pengaruh sentimen pada suatu kalimat. Proses *stopword* yang digunakan pada penelitian ini memanfaatkan *library* Sastrawi yang terdapat *corpus stopwords* bahasa Indonesia. Tabel 6 Berikut merupakan hasil *stopword removal* pada *data set* dalam penelitian ini.

Tabel 6 Hasil *Stopword Removal*

Sebelum	Sesudah
['sangat', 'bagus', 'dan', 'membantu']	['bagus', 'membantu']

6). Stemming

Stemming merupakan tahap untuk menghilangkan imbuhan pada suatu kata sehingga menjadi kata dasar dengan tujuan untuk menyamakan kata-kata yang ada pada dokumen *data set* agar menjadi kata dasar. *Stemming* pada *data set* berbahasa Indonesia dilakukan dengan memanfaatkan *library* Sastrawi pada bahasa pemrograman *Python*. Tabel 7 adalah contoh hasil stemming pada *data set*.

Tabel 7 Hasil *Stemming*

Sebelum	Sesudah
['bagus', 'membantu']	['bagus', 'bantu']

B. Pelabelan Data

Data set pada penelitian ini diberi label berdasarkan *rating* pada *Google Play* yang dibagi menjadi tiga kelas label yaitu positif, netral, dan negatif. *Review* dengan *rating* 4 dan 5 diberikan label positif, *rating* 3 berlabel netral, sementara itu *rating* 2 dan 1 diberikan label negatif.

Tabel 8 Hasil Pelabelan

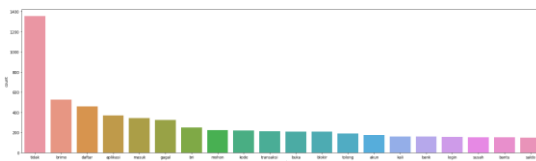
Label	Jumlah
Positif	48.178
Netral	2.179
Negatif	13.288

Tabel 9 Merupakan hasil pelabelan 63.644 *data set* bersih sehingga dapat diketahui bahwa 48.178 *data set* data dengan label positif, 2.179 data dengan label netral, dan 13.288 data dengan label negatif. Berdasarkan hasil pemberian label berdasarkan, diketahui bahwa sentimen pada *review* aplikasi mobile banking bank x dari tanggal 1 Januari – 27 Juni 2022 adalah positif.

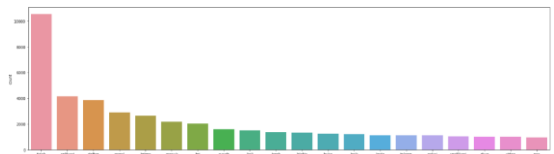
Year	Number of Publications
1990	100
1991	120
1992	140
1993	160
1994	180
1995	200
1996	220
1997	240
1998	260
1999	280
2000	300
2001	320
2002	340
2003	360
2004	380
2005	400
2006	420
2007	440
2008	460
2009	480
2010	500
2011	520
2012	540
2013	560
2014	580
2015	600
2016	620
2017	640
2018	660
2019	680

[illegible]

Gambar 5 adalah grafik kemunculan kata dalam kelas netral, sementara Gambar 6 ialah visualisasi *wordcloud* dari kelas netral dengan kata tidak, masuk, daftar, gagal, mohon, tolong, susah, bantu merupakan kata yang paling sering muncul.

[illegible]

Gambar 7 adalah grafik kemunculan kata dalam kelas negatif, sementara Gambar 8 ialah visualisasi *wordcloud* dari kelas negatif dengan kata tidak, daftar, susah, blokir, tolong, verifikasi merupakan kata yang paling sering muncul.

[illegible]

229

C. Pembobotan TF-IDF

Pada tahap ini, dilakukan transformasi *data set* yang masih berbentuk teks menjadi bentuk numerik atau angka menggunakan metode pembobotan TF-IDF. Tabel 10 merupakan hasil dari pembobotan dengan metode TF-IDF.

Tabel 9 Hasil TF-IDF

Term	Weight
mudah	0.073549
bantu	0.056695
transaksi	0.052267
brimo	0.042184
aplikasi	0.039928

3. Eksperimen dan Pengujian Metode

A. Data Latih dan Data Uji

Tahap selanjutnya adalah membagi *data set* menjadi data latih dan data uji. Data latih digunakan untuk melatih algoritma klasifikasi berdasarkan *data set* dalam penelitian ini. Data uji merupakan data yang digunakan untuk menguji kinerja dari algoritma klasifikasi dimana kinerja tersebut dihitung berdasarkan data benar yang diklasifikasi. Tabel 11 adalah perbandingan data latih dan data uji yang digunakan dalam penelitian.

Tabel 10 Perbandingan Data Latih dan Data Uji

Data Latih: Data Uji	Data latih	Data Uji
90% : 10%	57.280	6.365
80% : 20%	50.916	12.729
70% : 30%	4551	19.094

B. Klasifikasi Naïve Bayes

Setelah pembagian *data set* menjadi data latih dan data uji, dilakukan implementasi algoritma untuk klasifikasi data uji berdasarkan data latih. Tahap ini dilakukan tiga kali dengan masing-masing rasio data latih dan data uji. Hasil akurasi dengan pembagian *data set* 90% : 10% adalah 87,82% untuk pembagian data 80% : 20% akurasi adalah 86,87% dan untuk pembagian data 70% : 30% menghasilkan akurasi sebesar 86,88%. Tingkat akurasi terbaik pada algoritma *Naïve Bayes* yaitu pada rasio data latih 90% : 10% data uji. Tabel 12 adalah kinerja algoritma *Naïve Bayes* dengan perbedaan data latih dan data uji.

Tabel 11 Klasifikasi *Naïve Bayes*

Data Latih : Data Uji	Akurasi	Presisi	Recall
90% : 10%	87,82%	58,90%	60,08%
80% : 20%	86,87%	58,66%	60,30%
70% : 30%	86,88%	58,66%	60,30%

Hasil *confusion matrix* pada metode *Naïve Bayes* dengan rasio 90% : 10% didapatkan bahwa prediksi benar pada sentimen positif atau *true positive* sebanyak 4482, prediksi benar pada sentimen netral atau *true netral* sebanyak 12, dan prediksi benar pada sentimen negatif atau *true negative* sebanyak 1096. Tabel 13 adalah *confusion matrix* pada metode *Naïve Bayes*.

Tabel 12 Hasil *Confusion matrix* Metode *Naïve Bayes*

Hasil Aktual	Nilai Prediksi		
	Negatif	Netral	Positif
Negatif	1096	76	176
Netral	124	12	76
Positif	266	57	4482

Persamaan 1 adalah perhitungan nilai akurasi dari metode *Naïve Bayes*

$$\text{Akurasi} = \frac{\text{total prediksi benar}}{\text{total keseluruhan data}} = \frac{5590}{6365} = 87,82\% \quad (1)$$

Persamaan 2, 3, dan 4 adalah perhitungan presisi kelas positif, netral dan negatif dari metode *Naïve Bayes*

$$\begin{aligned} \text{Presisi kelas positif} &= \frac{\text{true positif}}{\text{total prediksi positif}} \\ &= \frac{4482}{4482+76+176} = 94,68\% \quad (2) \end{aligned}$$

$$\begin{aligned} \text{Presisi kelas netral} &= \frac{\text{true netral}}{\text{total prediksi netral}} \\ &= \frac{12}{76+12+57} = 8,28\% \quad (3) \end{aligned}$$

$$\begin{aligned} \text{Presisi kelas negatif} &= \frac{\text{true negatif}}{\text{total prediksi negatif}} \\ &= \frac{1096}{1096+124+266} = 73,76\% \quad (4) \end{aligned}$$

Persamaan 5, 6, dan 7 adalah perhitungan *recall* kelas positif, netral dan negatif dari metode *Naïve Bayes*

$$\begin{aligned} \text{Recall kelas positif} &= \frac{\text{true positif}}{\text{data positif aktual}} \\ &= \frac{4482}{4482+57+266} = 93,28\% \quad (5) \end{aligned}$$

$$\begin{aligned} \text{Recall kelas netral} &= \frac{\text{true netral}}{\text{data netral aktual}} \\ &= \frac{12}{124+12+76} = 5,66\% \quad (6) \end{aligned}$$

$$\begin{aligned} \text{Recall kelas negatif} &= \frac{\text{true negatif}}{\text{data negatif aktual}} \\ &= \frac{1096}{1096+76+176} = 81,31\% \quad (7) \end{aligned}$$

Setelah rasio dengan nilai akurasi terbaik diketahui, tahap selanjutnya dilakukan *cross validation* untuk mendapatkan nilai akurasi, presisi, dan *recall* yang maksimal. Pada penelitian ini digunakan *K-folds cross validation* dengan nilai $K=10$, dan mendapatkan

hasil akurasi terbaik pada iterasi ke-3 sejumlah 94%, yang memiliki nilai presisi serta *recall* sebesar 48% dan 58%.

Tabel 14 adalah hasil dari *10-folds cross validation*.

Tabel 13 Hasil *10-Folds Cross Validation* Metode *Naïve Bayes*

n	Akurasi	Presisi	Recall
1	68%	50%	51%
2	85%	58%	59%
3	94%	48%	58%
4	93%	49%	60%
5	92%	52%	58%
6	86%	56%	60%
7	87%	59%	60%
8	89%	60%	61%
9	89%	59%	60%
10	87%	59%	60%

C. Klasifikasi Support Vector Machine

Dilakukan implementasi algoritma *Support Vector Machine* untuk klasifikasi pada data uji berdasarkan data latih. Tahap ini dilakukan tiga kali dengan rasio data latih dan data uji yang berbeda. Hasil akurasi dengan pembagian *data set* 90% : 10% adalah 90,01% untuk pembagian data 80% : 20% akurasinya adalah 89,56% dan untuk pembagian data 70% : 30% menghasilkan akurasi sebesar 89,35%. Tingkat Akurasi terbesar pada algoritma *Support Vector Machine* yaitu pada rasio data latih 90% : 10% data uji. Tabel 15 adalah kinerja algoritma *Support Vector Machine* dengan perbedaan data latih dan data uji.

Tabel 14 Klasifikasi *Support Vector Machine*

Data Latih : Data Uji	Akurasi	Presisi	Recall
90% : 10%	90,01%	61,89%	60, 18%
80% : 20%	89,56%	59,82%	59,68%
70% : 30%	89,35%	66,10%	59,78%

Hasil *confusion matrix* pada metode *Support Vector Machine* dengan rasio 90% : 10% didapatkan bahwa bahwa prediksi benar pada sentimen positif atau *true positive* sebanyak 4588, prediksi benar pada sentimen netral atau *true netral* sebanyak 1, dan prediksi benar pada sentimen negatif atau *true negative* sebanyak 1140. Tabel 16 adalah *confusion matrix* pada metode *Support Vector Machine*.

Tabel 15 Hasil *Confusion matrix* Metode *Support Vector Machine*

Hasil Aktual	Nilai Prediksi		
	Negatif	Netral	Positif
Negatif	1140	2	206
Netral	118	1	93
Positif	213	4	4588

Persamaan 8 adalah perhitungan nilai akurasi dari metode *Support Vector Machine*.

$$\text{Akurasi} = \frac{\text{total prediksi benar}}{\text{total keseluruhan data}} = \frac{5729}{6365} = 90,01\% \quad (8)$$

Persamaan 9, 10, dan 11 adalah perhitungan presisi kelas positif, netral dan negatif dari metode *Support Vector Machine*

$$\begin{aligned} \text{Presisi kelas positif} &= \frac{\text{true positif}}{\text{total prediksi positif}} \\ &= \frac{4588}{4588+93+206} = 93,88\% \quad (9) \end{aligned}$$

$$\begin{aligned} \text{Presisi kelas netral} &= \frac{\text{true netral}}{\text{total prediksi netral}} \\ &= \frac{1}{2+1+4} = 14,29\% \quad (10) \end{aligned}$$

$$\begin{aligned} \text{Presisi kelas negatif} &= \frac{\text{true negatif}}{\text{total prediksi negatif}} \\ &= \frac{1140}{1140+118+213} = 77,50\% \quad (11) \end{aligned}$$

Persamaan 12, 13, dan 14 adalah perhitungan *recall* kelas positif, netral dan negatif dari metode *Support Vector Machine*

$$\begin{aligned} \text{Recall kelas positif} &= \frac{\text{true positif}}{\text{data positif aktual}} \\ &= \frac{4588}{4588+4+213} = 93,28\% \quad (12) \end{aligned}$$

$$\begin{aligned} \text{Recall kelas netral} &= \frac{\text{true netral}}{\text{data netral aktual}} \\ &= \frac{1}{118+1+93} = 0,47\% \quad (13) \end{aligned}$$

$$\begin{aligned} \text{Recall kelas negatif} &= \frac{\text{true negatif}}{\text{data negatif aktual}} \\ &= \frac{1140}{1140+2+206} = 84,57\% \quad (14) \end{aligned}$$

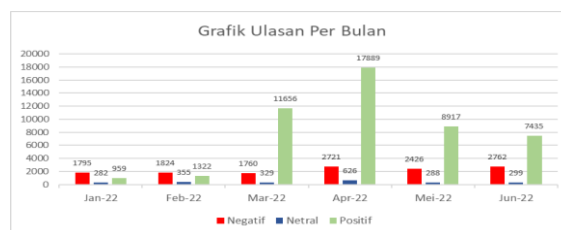
Setelah rasio dengan nilai akurasi terbaik diketahui, tahap selanjutnya dilakukan *cross validation* untuk mendapatkan nilai akurasi, presisi, dan *recall* yang maksimal. Pada penelitian ini digunakan *K-folds cross validation* dengan nilai $K=10$, dan mendapatkan hasil akurasi terbaik pada iterasi ke-3 sejumlah 96%, dan memiliki nilai presisi erta *recall* sebesar 50% dan 58%. Tabel 17 adalah hasil dari *10-folds cross validation*.

Tabel 16 Hasil *10-Folds Cross Validation* Metode *Support Vector Machine*

n	Akurasi	Presisi	Recall
1	71%	53%	53%
2	87%	64%	59%
3	96%	50%	58%
4	95%	65%	59%
5	94%	51%	59%
6	88%	70%	60%
7	89%	57%	59%
8	92%	59%	61%
9	91%	58%	60%
10	90%	58%	60%

4. Interpretasi Hasil

Penelitian ini melakukan klasifikasi sentimen data ulasan *Google Play* dengan menggunakan algoritma *Naïve Bayes* dan *Support Vector Machine* dengan tiga rasio pembagian data latih dan data uji yang berbeda. Pengambilan data menggunakan teknik *scraping* menghasilkan 8996 data mentah. Kemudian dilakukan *text preprocessing* menjadi 63.645 data bersih. Setelah itu dilakukan pelabelan 63.645 *data set* menggunakan *rating* dengan tiga kelas label, dimana *rating* 1 dan 2 berlabel negatif, 3 berlabel netral, sedangkan *rating* 4 dan 5 berlabel positif. Berdasarkan hasil pelabelan tersebut didapatkan total sebanyak 48.178 ulasan positif, 2.179 ulasan netral, dan 13.288 ulasan negatif. Gambar 9 merupakan grafik ulasan masing-masing kelas per bulan, dimana bulan April memiliki jumlah ulasan positif, netral, dan negatif terbanyak yaitu sejumlah 17.889 ulasan positif, 626 ulasan netral, dan 2721 ulasan negatif.



Gambar 9 Grafik Ulasan Tiap Kelas Per Bulan

Berdasarkan proses pengujian klasifikasi menggunakan metode *Naïve Bayes* dan *Support Vector Machine* pada data ulasan didapatkan hasil akurasi dari masing-masing metode tersebut pada Tabel 18. Hasil dari penelitian klasifikasi teks data ulasan aplikasi *mobile banking* pada situs *Google Play* menggunakan metode *Naïve Bayes* mendapatkan akurasi 87,82%, sedangkan metode *Support Vector Machine* mendapatkan akurasi 90,01% dengan rasio pembagian data latih dan data uji 90%:10%.

Tabel 17 Kinerja *Naïve Bayes* dan *Support Vector Machine*

Rasio	<i>Naïve Bayes</i>			<i>Support Vector Machine</i>		
	Akurasi	Presisi	Recall	Akurasi	Presisi	Recall
90% : 10%	87,82 %	58,9 0%	60,0 8%	90,01 %	61,8 9%	60, 18%
80% : 20%	86,87 %	58,6 6%	60,3 0%	89,56 %	59,8 2%	59,6 8%
70% : 30%	86,88 %	58,6 6%	60,3 0%	89,35 %	66,1 0%	59,7 8%

Pada penelitian T. Kurniawan (Kurniawan, 2017) dalam analisis sentimen berjumlah 1.000 *data set* dengan metode *Naïve Bayes* dan *Support Vector Machine* mendapatkan tingkat akurasi 97,8% dan 99,3%. Penelitian Ilmawan dan Mude (Ilmawan & Mude, 2020) menggunakan metode *Naïve Bayes* dan *Support Vector Machine* dengan 1818 komentar *review*, menghasilkan akurasi sebesar 75,41% dan 81,46%.

Penelitian lain oleh Ahmad (Ahmad, 2020) untuk analisis sentimen menggunakan metode *Naïve Bayes* dan *Support Vector Machine* berjumlah 9015 data tweet terkait COVID-19 dengan kata kunci “lockdown, psbb, karantina” menghasilkan akurasi 82,3% dan 83%.

Penelitian ini dilakukan pengujian 8996 ulasan aplikasi *mobile banking* pada situs *Google Play*. *Data set* yang berhasil di-*scraping* merupakan data mentah yang tidak terstruktur sehingga perlu dilakukan pembersihan atau tahap *preprocessing* agar lebih mudah dikenali bentuknya oleh sistem. Tahap selanjutnya adalah proses pelabelan dimana *data set* dibagi menjadi kelas positif, netral, dan negatif berdasarkan *rating* ulasan pengguna. Implementasi metode klasifikasi diterapkan pada *data set* dimana hasil akurasi maksimal metode *Naïve Bayes* adalah 87,82% dan akurasi maksimal metode *Support Vector Machine* adalah 90,01%. Perbandingan kinerja antara penelitian ini dengan penelitian lain dapat dilihat pada Tabel 19.

Tabel 18 Perbandingan kinerja NB dan SVM Terhadap Penelitian Lain

Peneliti	Metode	
	NB	SVM
(Kurniawan, 2017)	97,8%	99,3%
(Ilmawan & Mude, 2020)	75,41%	81,46%
(Ahmad, 2020)	82,3%	83%
(Longdong, 2022)	87,82%	90,01%

PENUTUP

1. KESIMPULAN

Berdasarkan pembahasan diatas maka penulis menarik kesimpulan sebagai berikut:

- Dari hasil pelabelan berdasarkan *rating* pada *data set* bersih sejumlah 63.645 data ulasan, didapati total sebanyak 48.178 ulasan positif, 2.179 ulasan netral, dan 13.288 ulasan negatif. Hal ini menggambarkan bahwa sentimen pengguna aplikasi *mobile banking* bank X secara umum dominan positif.
- Implementasi metode *Naïve Bayes* dan *Support Vector Machine* pada penelitian ini mendapatkan hasil akurasi maksimal pada metode *Naïve Bayes* dan *Support Vector Machine* sebesar 87,82% dan 90,01% seperti tertuang pada Tabel 18. Berdasarkan hasil akurasi tersebut dapat diketahui bahwa metode *Support Vector Machine* memiliki tingkat performansi yang lebih tinggi dari metode *Naïve Bayes* untuk analisis sentimen data ulasan aplikasi *mobile banking* bank X, pada penelitian ini.

2. SARAN

Saran yang penulis berikan untuk penelitian selanjutnya antara lain:

- Pada penelitian ini didapatkan bahwa metode *Support Vector Machine* memiliki kinerja yang lebih baik sehingga

pada penelitian berikutnya dapat menggunakan metode tersebut untuk melakukan analisis sentimen.

- b. Penelitian ini merupakan tipe fine-grained sentimen analisis dimana sentimen didasari oleh *rating* dan teks ulasan pengguna terhadap suatu aplikasi. Pada penelitian berikutnya dapat

menggunakan tipe analisis sentimen dengan *emotion detection* untuk menganalisis emosi yang terkandung dalam sebuah teks ulasan atau menggunakan *aspect base* untuk mengetahui aspek apa yang menjadi penilaian pengguna terhadap suatu aplikasi.

DAFTAR PUSTAKA

- Ahmad, F. (2020). *Studi Perbandingan Metode Analisis Naive Bayes Classifier Dengan Support Vector Machine Untuk Analisis Sentimen (Studi Kasus: Tweet Berbahasa Indonesia Tentang Covid-19)*. 152.
- Anjasmoros, M. T., Marisa, F., & Istiadi. (2020). Analisis Sentimen Aplikasi Go-Jek Menggunakan Metode Svm Dan Nbc (Studi Kasus: Komentar Pada Play Store). *Conference on Innovation and Application of Science and Technology (CIASTECH 2020)*, (Ciastech), 489–498.
- Awaludin, M. (2015). Penerapan Metode Distance Transform Pada Linear Discriminant Analysis Untuk Kemunculan Kulit Pada Deteksi Kulit. *Journal of Intelligent Systems*, 1(1), 49–55.
- Awaludin, M., & Ridyustia Raveena, R. (2021). Penerapan Metode Rational Unified Process Pada Knowledge Management System Untuk Mendukung Proses Pembelajaran Sekolah Menengah Atas. *JSI (Jurnal Sistem Informasi) Universitas Suryadarma*, 8(2), 159–170.
- Awaludin, M., Yasin, V., & Risyda, F. (2024). The Influence of Artificial Intelligence Technology, Infrastructure and Human Resource Competence on Internet Access Networks. *Inform : Jurnal Ilmiah Bidang Teknologi Informasi Dan Komunikasi*, 9(2), 111–120. <https://doi.org/10.25139/inform.v9i2.8109>
- Bahtera, E. G., Vidyarini, T. N., & ... (2019). Citra Shopee Pasca Kasus Petisi Pemboikotan Iklan Versi Blackpink “12.12 Birthday Sale” di Media Online. *Jurnal E-Komunikasi*.
- Hussein, S. (2021). Support Vector Machine, Algoritma untuk Machine Learning – GEOSPASIALIS.
- Ilmawan, L. B., & Mude, M. A. (2020). Perbandingan Metode Klasifikasi Support Vector Machine dan Naïve Bayes untuk Analisis Sentimen pada Ulasan Tekstual di Google Play Store. *ILKOM Jurnal Ilmiah*, 12(2), 154–161. <https://doi.org/10.33096/ilkom.v12i2.597.154-161>
- Kurniawan, T. (2017). Implementasi Text Mining Pada Analisis Sentimen Pengguna Twitter Terhadap Media Mainstream Menggunakan Naïve Bayes Classifier Dan Support Vector Machine Media Mainstream. *IT Journal*, 23, 1.

- Litbang, A. S. (2021). Mengenal Metode Analisis Klasifikasi Naive Bayes.
- Masturoh, S. (2021). Analisis Sentimen Terhadap E-Wallet Dana Pada Ulasan Google Play Menggunakan Algoritma K-Nearest Neighbor. *Jurnal Pilar Nusa Mandiri*, 17(1), 53–58. <https://doi.org/10.33480/pilar.v17i1.2182>
- Owen, L. (2020). NLP Bahasa Indonesia Resources.
- Python Software Foundation. (2021). google-play-scraper · PyPI.