

IMPLEMENTASI ALGORITMA MULTINOMIAL NAIVE BAYES PADA ANALISIS SENTIMEN TERHADAP ULASAN PENGGUNA APLIKASI BRIMO

Nanda Fibriyanti Arminda, Nina Sulistiyowati, Tesa Nur Padilah

Program Studi Informatika S1, Fakultas Ilmu Komputer, Universitas Singaperbangsa Karawang
Jl.HS. Ronggo Waluyo, Telukjambe Timur, Karawang, Indonesia
1910631170219@student.unsika.ac.id

ABSTRAK

Mobile banking merupakan salah satu produk perbankan yang dapat memberikan kemudahan bagi nasabah. Salah satu *mobile banking* populer adalah BRImo. Ulasan pengguna menjadi sumber informasi penting bagi developer untuk memahami keluhan pengguna serta meningkatkan pengembangan aplikasi. Namun tidak semua ulasan mewakili pendapat tentang aplikasi tersebut. Penelitian ini bertujuan untuk melakukan analisis sentimen guna mengekstrak informasi dan mengklasifikasikan data ulasan ke dalam sentimen positif dan negatif dengan mengimplementasikan algoritma Naïve Bayes dalam proses klasifikasi dan menggunakan *K-Fold Cross Validation* sebagai metode pembagian data dan metode validasi hasil evaluasi, serta digunakan teknik TF-IDF untuk melakukan transformasi data dan pembobotan kata. Metode penelitian dilakukan dengan menerapkan tahapan metode KDD (*Knowledge Discovery in Database*) dengan tahapan *data cleaning*, *data selection*, *data transformation*, *data mining*, *pattern evaluation*, dan *knowledge presentation*. Hasil analisis sentimen ulasan pengguna BRImo dari total *dataset* yang berjumlah 1011 data terdapat 670 data yang diklasifikasikan sebagai sentimen negatif dan 341 data yang diklasifikasikan sebagai sentimen positif, sehingga dapat disimpulkan bahwa hasil sentimen dari ulasan pengguna terhadap aplikasi BRImo didominasi oleh sentimen negatif, dengan hasil evaluasi dari algoritma Naïve bayes dalam mengklasifikasikan sentimen ulasan pengguna BRImo mendapatkan hasil terbaik pada *fold-2* dengan nilai akurasi 98,02%, presisi 97,06%, *recall* 97,06%, dan *f1-score* 97,06%.

Kata kunci: Analisis Sentimen, BRImo, K-Fold Cross Validation, KDD, Multinomial Naïve Bayes, TF-IDF.

1. PENDAHULUAN

Dalam perkembangan zaman yang semakin pesat dan kompetitif ini transformasi digital pada layanan perbankan sudah menjadi hal yang umum, hal tersebut dilakukan guna mempertahankan dan memperoleh nasabah. Dengan adanya layanan perbankan digital, pihak bank melakukan orientasi penuh untuk kebutuhan nasabahnya dengan menggunakan teknologi digital. Selain itu, layanan perbankan digital juga dapat memberikan kemudahan bagi nasabah karena transaksi dapat dengan mudah dilakukan hanya melalui telepon seluler [1].

Salah satu jenis layanan perbankan digital yang saat ini banyak digunakan yaitu *mobile banking*. Moda transaksi *mobile banking* saat ini sudah menjadi hal yang populer dan lumrah di kalangan masyarakat [2]. Salah satu aplikasi *mobile banking* yang banyak digunakan di kalangan masyarakat adalah BRImo. Hal ini dibuktikan berdasarkan jumlah unduhan dari aplikasi BRImo yang telah mencapai lebih dari 10 juta, dengan total ulasan mencapai lebih dari 1 juta pada *Google Play Store*. Hal menarik dari *Google Play Store* ini yaitu terdapat fitur yang berisi ulasan dari para pengguna untuk suatu aplikasi tertentu. Umumnya ulasan tersebut dapat digunakan sebagai tolak ukur untuk menemukan informasi yang efektif dan efisien dari suatu produk [3].

Ulasan pengguna dari suatu aplikasi seperti halnya BRImo biasanya berisi saran positif juga keluhan yang sifatnya negatif, hal ini tentunya akan mempengaruhi ekspektasi dari calon pengguna. Maka,

diperlukan metode yang secara cepat dan otomatis dapat menyortir dan mengategorikan ulasan tersebut. *Text mining* merupakan suatu proses ekstraksi informasi berkualitas tinggi dari data teks untuk mengidentifikasi masalah yang terkait dengan topik tertentu. Dalam analisis sentimen, *text mining* dapat mengenali sentimen yang terkandung dalam pernyataan tersebut [4]. Selain itu, besarnya manfaat dan pengaruh analisis sentimen mengakibatkan perkembangan penelitian maupun aplikasi mengenai analisis sentimen semakin pesat [5].

Adapun algoritma pengklasifikasian dalam penelitian ini akan memanfaatkan algoritma *Naïve Bayes Classifier*. Pemilihan algoritma ini didasarkan pada penelitian sebelumnya yang dilakukan oleh Ting, Ip, dan Tsang (2011) yang berjudul "*Is Naïve Bayes a Good Classifier for Document Classification?*" pada penelitian ini penggunaan algoritma Naïve Bayes dianggap berpotensi baik dalam mengklasifikasikan data jika dibandingkan dengan metode klasifikasi lainnya (*Neural Network*, *SVM*, dan *Decision Tree*) dalam hal akurasi dan komputasi. Berdasarkan hasil evaluasi yang didapatkan pada penelitian ini, nilai evaluasi dari Naïve Bayes merupakan yang terbaik. Meskipun *SVM* mendapatkan hasil yang serupa dengan Naïve Bayes, namun waktu yang dibutuhkan untuk membuat model pada penelitian ini tidak memuaskan [6]. Algoritma Naïve Bayes juga telah banyak digunakan karena kesederhanaannya baik dalam tahap pelatihan maupun pengklasifikasian. Selain itu, Naïve Bayes memungkinkan setiap atribut

untuk berkontribusi terhadap keputusan akhir secara sama dan independen dari atribut lainnya, sehingga lebih efisien dalam hal komputasi bila dibandingkan dengan pengklasifikasi teks lainnya [4].

Berdasarkan penelitian yang telah dilakukan sebelumnya, maka pada penelitian ini akan dilakukan analisis sentimen ulasan pengguna terhadap kinerja dari aplikasi BRImo pada *Google Play Store* dengan menggunakan metode *Multinomial Naïve Bayes*.

2. TINJAUAN PUSTAKA

Pada bagian ini dipaparkan literatur ilmiah yang relevan dengan penelitian yang dilakukan saat ini.

2.1. Analisis Sentimen

Analisis sentimen adalah bidang studi komputasi yang mengkaji opini, perilaku, sentimen, evaluasi, serta emosi seseorang terhadap berbagai entitas. Adapun entitas tersebut dapat berupa gambaran tentang produk, peristiwa, layanan, organisasi, masalah, individu, atau topik tertentu [3]. Tujuan dari analisis sentimen adalah untuk mengidentifikasi atribut dan komponen dari entitas yang diulas dalam sebuah dokumen, dengan maksud menentukan apakah komentar tersebut mengungkapkan sentimen positif atau negatif. [5].

2.2. Term Frequency–Inverse Document Frequency (TF–IDF)

TF–IDF merupakan teknik pembobotan yang biasanya digunakan untuk mentransformasikan data teks ke dalam format numerik. TF-IDF sendiri merupakan perpaduan dari algoritma pembobotan berbeda, yaitu TF yang dapat menghitung jumlah kemunculan kata dalam suatu dokumen, dan IDF yang menunjukkan tingkat pentingnya suatu kata dalam kumpulan dokumen. Dalam TF-IDF, fitur atau *term* penting akan memiliki nilai bobot yang tinggi, sedangkan *term* yang kurang penting akan memiliki nilai bobot rendah [7].

2.3. K-fold Cross Validation

K-fold Cross Validation merupakan salah satu teknik *sampling* yang digunakan untuk membagi *dataset* ke dalam beberapa bagian yang biasanya digunakan untuk menilai atau memvalidasi tingkat akurasi suatu model yang dibangun berdasarkan *dataset* tertentu dengan meminimalkan waktu komputasi dan tetap menjaga akurasi dari estimasi. Cara kerja dari metode ini yaitu dengan membagi secara acak data ke dalam *k* kelompok dari perkiraan pembagian nilai ukuran yang setara seperti pada *dataset* yang penuh [8].

2.4. Multinomial Naïve Bayes

Multinomial Naïve Bayes merupakan suatu pendekatan yang bekerja dengan menghitung frekuensi dari masing-masing *term* dalam dokumen [9]. Kategori dokumen pada klasifikasi *Multinomial Naïve Bayes* tidak hanya berdasarkan kata-kata yang

muncul di dalamnya, tetapi juga mempertimbangkan jumlah kemunculan dari kata-kata tersebut. Sehingga dalam konteks *Multinomial Naïve Bayes* ini peran tokenisasi sangat penting [10].

2.5. Confusion Matrix

Confusion matrix adalah sebuah metode dalam data mining yang biasanya digunakan untuk mengevaluasi performa suatu algoritma atau model klasifikasi, dengan menunjukkan keakuratan, presisi, dan tingkat kesalahan algoritma tersebut. *Confusion matrix* biasanya direpresentasikan dalam bentuk tabel yang menunjukkan empat kemungkinan hasil klasifikasi dengan matriks baris merupakan jumlah data uji untuk kelas sebenarnya sedangkan pada bagian kolom merupakan kelas prediksi, seperti berikut [3].

Tabel 1. *Confusion Matrix*

		Kelas Sebenarnya	
		True	False
Kelas Prediksi	True	TP	FP
	False	FN	TN

Keterangan:

- 1) *True Positive* (TP), merupakan jumlah data yang dengan benar diklasifikasikan sebagai kelas positif.
- 2) *True Negative* (TN), merupakan jumlah data yang benar diklasifikasikan sebagai kelas negatif.
- 3) *False Positive* (FP), merupakan jumlah data yang seharusnya adalah kelas negatif, namun salah diklasifikasikan sebagai kelas positif.
- 4) *False negative* (FN), merupakan jumlah data yang seharusnya adalah kelas positif, namun salah diklasifikasikan sebagai kelas negatif.

Pengujian dengan menggunakan *confusion matrix* menghasilkan perhitungan yang menghasilkan empat keluaran, antara lain yaitu:

- 1) Akurasi, digunakan untuk mengukur ketepatan suatu model dalam melakukan klasifikasi data secara keseluruhan, berikut rumus menghitung nilai akurasi.

$$Accuracy = \frac{TP+TN}{(TP+TN+FP+FN)} \quad (1)$$

- 2) Presisi, digunakan untuk mengukur banyaknya prediksi positif yang benar terhadap total prediksi positif, serta meminimalkan *false positive*. berikut rumus menghitung nilai presisi.

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

- 3) *Recall*, digunakan untuk mengukur banyaknya prediksi positif yang benar terhadap total data aktual yang positif, serta meminimalkan *false negative*, berikut rumus menghitung nilai *recall*.

$$Recall = \frac{TP}{TP+FN} \quad (3)$$

- 4) *F1 Score*, adalah *harmonic mean* dari presisi dan *recall* yang berguna untuk mengevaluasi model dengan presisi dan *recall* yang tinggi, berikut rumus menghitung nilai *F1 score*.

$$F1\ Score = 2 \frac{(precision \times recall)}{precision + recall} \quad (4)$$

Keterangan :

TP : Jumlah prediksi *true positive*

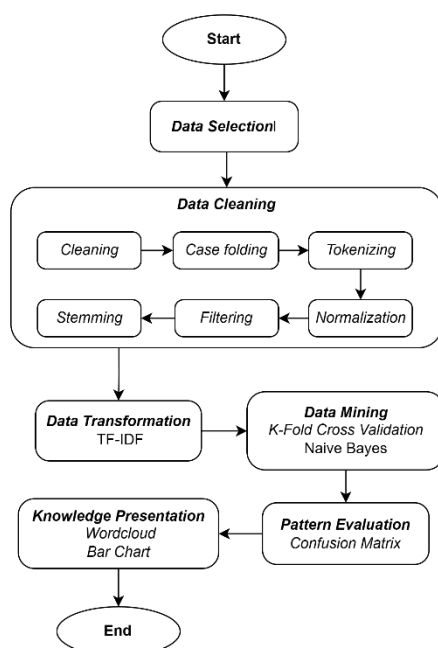
TN : Jumlah prediksi *true negative*

FP : Jumlah prediksi *false positive*

FN : Jumlah prediksi *false negative*

3. METODE PENELITIAN

Metodologi penelitian yang akan diterapkan pada penelitian ini yaitu metode *Knowledge Discovery in Data* (KDD) yang terdiri dari tahapan *data selection*, *data cleaning*, *data transformation*, *data mining*, *pattern evaluation*, dan *knowledge presentation*. Berikut merupakan gambaran mengenai perincian metode penelitian yang akan dilaksanakan.



Gambar 1. Metodologi Penelitian

3.1. Data Selection

Dalam tahap ini dilakukan proses *collecting data*, seleksi atribut, serta pemilihan data ulasan yang relevan dengan penelitian yang akan dilakukan. Selain itu, dilakukan pula pelabelan data secara manual dengan bantuan seorang ahli bahasa sebagai validator.

3.2. Data Cleaning

Pada tahap ini dilakukan pembuangan *noise*, data tidak konsisten, serta dilakukan perbaikan kesalahan penulisan dalam data dengan mengikuti tahapan dari

preprocessing data yaitu *cleaning*, *case folding*, *tokenizing*, *normalization*, *filtering*, dan *stemming*.

3.3. Data Transformation

Dalam tahap ini dilakukan transformasi data ke dalam format yang sesuai untuk dilakukan proses *data mining* nantinya, dengan menerapkan algoritma *Term Frequency-Inverse Document Frequency* (TF-IDF).

3.4. Data Mining

Pada tahap ini dilakukan proses pelacakan informasi maupun pola menarik menggunakan teknik *sampling K-Fold Cross Validation* dan algoritma *Multinomial naïve bayes* yang diaplikasikan pada data.

3.5. Pattern Evaluation

Dilakukan identifikasi model prediksi hasil dari proses sebelumnya yang telah dilakukan evaluasi menggunakan *confusion matrix* berupa hasil akurasi, presisi, *recall*, dan *f1-score* guna menilai capaian dari suatu hipotesis berdasarkan pola yang menarik.

3.6. Knowledge Presentation

Dilakukan visualisasi dan representasi pola informasi untuk membantu mengomunikasikan pengetahuan ke dalam bentuk yang mudah dimengerti. Pada penelitian ini informasi yang ditampilkan berupa kumpulan data yang divisualisasikan dalam bentuk *wordcloud*.

4. HASIL DAN PEMBAHASAN

Penelitian ini dilakukan dengan menggunakan bahasa pemrograman *python* dengan *tools google colab*, dan mengimplementasikan algoritma *Multinomial Naïve Bayes Classifier* untuk membuat model analisis sentimen terhadap ulasan pengguna aplikasi BRImo pada *Google Play Store* dengan menggunakan metode penelitian *Knowledge Discovery in Database* (KDD) yang diimplementasikan dalam enam tahapan sebagai berikut.

4.1. Data Selection

Dalam proses ini dilakukan *crawling* data ulasan pengguna aplikasi BRImo berbahasa Indonesia pada *Google Play Store* melalui *website AppFollow* dengan menggunakan kata kunci “BRImo” dan untuk periode pengambilan data dilakukan pada rentang waktu 28 Februari – 6 Maret 2023. Didapatkan 3620 data hasil *crawling* yang kemudian setelah dilakukan seleksi data didapatkan 1011 data yang akan digunakan pada penelitian ini, Data yang digunakan pada penelitian ini dapat dilihat pada Tabel 2 berikut.

Tabel 2. *Dataset* Penelitian

No	Review
1	sulit utk daftar, sdh dicoba berulang kali selalu gagal, bahkan di bri jg tdk berhasil.l
2	Memuaskan selalu aplikasi ini
3	Susah sekali untuk daftar harus ngerekam dan selalu gagal

No	Review
...	...
1010	Puas , nyaman dan mudah di operasikan
1011	Udah bikinnya susah.giliran udah jadi mau ligin malah salah usernime .udah verifikasi di coba lagi masih begitu terus ya ampun bikin susah aja

4.2. Data Cleaning

Proses pada *data cleaning* mengikuti tahapan dari *preprocessing data* sebagai berikut.

a. Cleaning

Dilakukan penghapusan atribut yang tidak berpengaruh terhadap klasifikasi. Seperti karakter simbol, angka, dan emotikon.

Tabel 3. *Cleaning*

Sebelum	Sesudah
Aplikasi yang Sangat Bermanfaat, Bgus BGT!!	Aplikasi yang Sangat Bermanfaat Bgus BGT

b. Case folding

Dilakukan konversi seluruh karakter dalam dokumen ke dalam format *lowercase*.

Tabel 4. *Case Folding*

Sebelum	Sesudah
Aplikasi yang Sangat Bermanfaat Bgus BGT	aplikasi yang sangat bermanfaat bgus bgt

c. Tokenizing

Dilakukan pemecahan kalimat ke dalam satuan kata (token) yang kemudian akan dipisahkan dengan karakter spasi atau koma.

Tabel 5. *Tokenizing*

Sebelum	Sesudah
aplikasi yang sangat bermanfaat bgus bgt	[aplikasi, yang, sangat, bermanfaat, bgus, bgt]

d. Normalization

Dilakukan pengoreksian kata yang mengandung unsur singkatan, tidak baku, dan typo.

Tabel 6. *Normalization*

Sebelum	Sesudah
[aplikasi, yang, sangat, bermanfaat, bgus, bgt]	[aplikasi, yang, sangat, bermanfaat, bagus, banget]

e. Filtering

Dilakukan penghapusan *stopwords* (kata-kata tidak deskriptif dan kurang penting).

Tabel 7. *Filtering*

Sebelum	Sesudah
[aplikasi, yang, sangat, bermanfaat, bagus, banget]	[aplikasi, bermanfaat, bagus, banget]

f. Stemming

Dilakukan transformasi kata dalam dokumen ke dalam bentuk kata dasar menggunakan aturan tertentu.

Tabel 8. *Stemming*

Sebelum	Sesudah
[aplikasi, bermanfaat, bagus, banget]	[aplikasi, manfaat, bagus, banget]

4.3. Data Transformation

Data yang telah diproses pada tahap sebelumnya diberi label untuk data positif dan negatif secara manual dan telah melalui proses validasi oleh ahli bahasa Indonesia untuk meminimalisasi kesalahan dalam proses pelabelan, kemudian *dataset* dan label yang sebelumnya berupa data tekstual diubah ke dalam format numerik dengan menggunakan algoritma TF-IDF (*Term Frequency-Inverse Document Frequency*) sebagai berikut.

	ada	adakan	admin	administrasi	adu	agen	agun	ahir	aja	ajar
0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
2	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
3	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
4	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
...	...	...	...	...	...	...	...	...	...	...
1006	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1007	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1008	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1009	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1010	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

1011 rows x 804 columns

Gambar 2. Matriks hasil TF-IDF

Dari gambar 2. diketahui bahwa hasil dari proses TF-IDF yang dilakukan berupa matriks yang berukuran 1011 x 804 dengan jumlah fitur sebanyak 804 *term* dan jumlah dokumen sebanyak 1011 data. Nilai pembobotan terhadap kemunculan kata dalam tiap ulasan tersebut kemudian digunakan untuk membantu proses klasifikasi pada tahap *data mining*.

4.4. Data Mining

Hasil dari *data mining* berupa implementasi dari algoritma *Multinomial Naïve Bayes* dengan proses pembagian *data training* dan *data testing* menggunakan teknik *sampling K-Fold Cross Validation* menggunakan jumlah K = 10. Hal ini dikarenakan nilai k=10 telah menjadi metode standar dalam hal praktis [11]serta merupakan nilai k yang paling umum digunakan karena memiliki kemampuan estimasi kinerja algoritma pembelajaran yang lebih akurat [12]. Kemudian hasil dari data yang telah melalui tahap ini dilakukan evaluasi menggunakan *confusion matrix* pada tahap *pattern evaluation*.

4.5. Pattern Evaluation

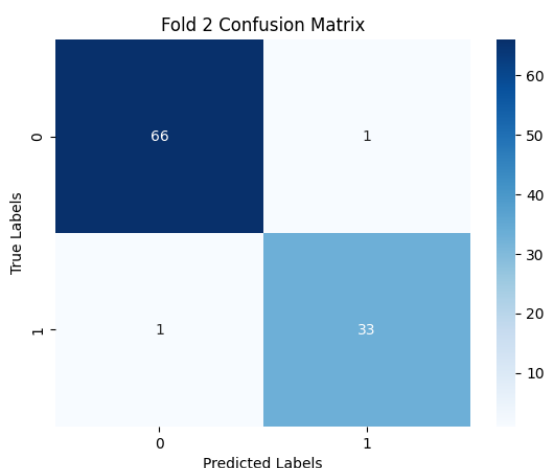
Hasil dari tahap ini berupa nilai evaluasi dari kinerja model dengan menggunakan metode evaluasi

confusion matrix yang disajikan dalam bentuk tabel berisi hasil evaluasi yang terdiri dari nilai akurasi, presisi, *recall*, *f1-score* dengan ukuran baris yang disesuaikan berdasarkan jumlah *k* yang telah diinisialisasikan pada tahap sebelumnya. Berikut merupakan hasil evaluasi dari model yang telah dibuat.

Tabel 9. Hasil Evaluasi

Fold	Akurasi	Presisi	Recall	F1-Score
ke-1	0,9314	0,9375	0,8571	0,8955
ke-2	0,9802	0,9706	0,9706	0,9706
ke-3	0,9406	0,9667	0,8529	0,9063
ke-4	0,9307	1,0	0,7941	0,8852
ke-5	0,9703	1,0	0,9118	0,9538
ke-6	0,9604	1,0	0,8824	0,9375
ke-7	0,9703	0,9697	0,9412	0,9552
ke-8	0,9703	0,9697	0,9412	0,9552
ke-9	0,9703	0,9697	0,9412	0,9552
ke-10	0,9703	0,9697	0,9412	0,9552

Berdasarkan tabel 9 dapat diketahui bahwa hasil pada *fold* ke-2 memiliki nilai evaluasi yang lebih baik dibandingkan dengan *fold* lainnya. Maka dari itu model klasifikasi sentimen pada *fold* ke-2 digunakan untuk evaluasi lebih lanjut dengan tabel *confusion matrix* sebagai berikut.



Gambar 3. Tabel *Confusion Matrix* Hasil Evaluasi *Fold-2*

Confusion matrix memberikan jumlah nilai prediksi benar (*true positive* dan *true negative*) dan jumlah nilai prediksi salah (*false positive* dan *false negative*) yang dapat digunakan untuk memperoleh hasil evaluasi seperti nilai akurasi, presisi, *recall* dan *f1-score*. Dapat diketahui dari gambar diatas bahwa jumlah data uji yang diprediksi benar adalah 99 data dan untuk data uji yang diprediksi salah adalah 2. Berikut merupakan perhitungan manual untuk memperoleh hasil evaluasi model.

$$\begin{aligned} 1) \text{ Akurasi} &= \frac{TP+TN}{(TP+TN+FP+FN)} = \frac{66+33}{(66+33+1+2)} = \frac{99}{101} \\ &= 0,9802 \approx \mathbf{0,98} \\ 2) \text{ Presisi} &= \frac{TP}{TP+FP} = \frac{33}{33+1} = \frac{33}{34} = 0,9706 \approx \mathbf{0,97} \\ 3) \text{ Recall} &= \frac{TP}{TP+FN} = \frac{33}{33+1} = \frac{33}{34} = 0,9706 \approx \mathbf{0,97} \end{aligned}$$

$$\begin{aligned} 4) \text{ F1 Score} &= 2 \frac{(\text{precision} \times \text{recall})}{(\text{precision} + \text{recall})} = 2 \frac{(0,9706 \times 0,9706)}{(0,9706 + 0,9706)} \\ &= 2 \frac{0,9420}{1,9412} = 2 \times 0,4853 = 0,9706 \\ &\approx \mathbf{0,97} \end{aligned}$$

4.6. Knowledge Presentation

Hasil dari tahap ini berupa visualisasi hasil penelitian untuk mengetahui gambaran atau informasi umum mengenai data ulasan pengguna aplikasi BRImo. Berikut merupakan *wordcloud* data ulasan positif terhadap aplikasi BRImo.



Gambar 4. *Wordcloud* Ulasan Positif Aplikasi BRImo



Gambar 5. *Wordcloud* Ulasan Negatif Aplikasi BRImo

Berdasarkan gambar 4. dapat dilihat bahwa semakin besar ukuran kata dalam *wordcloud* maka semakin tinggi frekuensi kemunculan kata tersebut dalam topik pembahasan pada ulasan positif. Adapun kata-kata yang sering muncul yaitu ‘mudah’, ‘aplikasi’, ‘bantu’, ‘transaksi’, ‘terima’, dan ‘kasih’. Hal ini menunjukkan bahwa ulasan positif dari pengguna aplikasi BRImo kerap berupa ungkapan kemudahan penggunaan aplikasi BRImo dalam membantu transaksi, dan ungkapan terima kasih terhadap kinerja aplikasi BRImo. Sedangkan untuk

wordcloud ulasan negatif terhadap BRImo ditampilkan pada gambar berikut.

Berdasarkan gambar 5. diketahui bahwa kata-kata yang sering muncul pada data ulasan negatif terhadap aplikasi BRImo yaitu ‘aplikasi’, ‘brimo’, ‘gagal’, ‘masuk’, ‘sulit’, dan ‘daftar’. Maka dapat disimpulkan bahwa pada ulasan negatif pengguna kerap menyampaikan bahwa kendala yang dialami ketika menggunakan aplikasi BRImo yaitu sering terjadi gagal masuk ke dalam aplikasi, dan sulitnya melakukan pendaftaran ketika akan melakukan registrasi akun BRImo. Beberapa topik tersebut dapat dimanfaatkan sebagai acuan bagi PT. Bank Rakyat Indonesia (Persero) Tbk. untuk evaluasi kinerja dari aplikasi BRImo demi kenyamanan pengguna.

5. KESIMPULAN DAN SARAN

Analisis sentimen pengguna terhadap aplikasi BRImo di *Google Play Store* yang terdiri dari total dataset berjumlah 1011 data, ditemukan jumlah ulasan sentimen negatif sebanyak 670 data sedangkan jumlah ulasan sentimen positif sebanyak 341 data. Hasil analisis sentimen pengguna terhadap aplikasi BRImo di *Google Play Store* didominasi oleh sentimen negatif yang diperoleh berdasarkan pemilihan kalimat secara manual yang dibantu oleh seorang ahli bahasa. Kata yang sering muncul pada ulasan negatif terhadap aplikasi BRImo adalah ‘aplikasi’, ‘brimo’, ‘gagal’, ‘masuk’, ‘sulit’, ‘daftar’, dan ‘rekening’. Maka dapat disimpulkan bahwa pengguna kerap mengeluhkan kinerja dari aplikasi BRImo terkait gagal masuk ke dalam aplikasi, sulitnya melakukan pendaftaran ketika akan melakukan registrasi akun, dan keluhan lainnya. Hasil evaluasi dari algoritma *Multinomial Naïve Bayes* yang telah dilakukan dalam mengklasifikasikan sentimen ulasan pengguna aplikasi BRImo menggunakan *10-fold cross validation* mendapatkan hasil evaluasi *confusion matrix terbaik* pada fold-2 dengan nilai akurasi 98,02%, presisi 97,06%, *recall* 97,06%, dan *f1-score* 97,06%. Untuk pengembangan lebih lanjut maka saran yang dapat diterapkan bagi penelitian selanjutnya yaitu, penggunaan algoritma lainnya untuk bisa menemukan nilai evaluasi yang lebih baik, seperti, *Support Vector Machine*, *Random Forest*, *K-Nearest Neighbor* dan lain sebagainya. Selain itu, dapat ditambahkan seleksi fitur seperti *Information Gain*, *Chi-Square*, *Particle Swarm Optimization*, atau seleksi fitur lainnya untuk mengoptimalkan parameter dengan memilih atribut yang relevan.

DAFTAR PUSTAKA

- [1] S. Mubarakah, “Minat menggunakan mobile banking pada perbankan milik negara,” in *NCAB*, Semarang, Oct. 2019, pp. 242–249.
- [2] E. Purwanto and J. Loisa, “The intention and use behaviour of the mobile banking system in Indonesia: UTAUT model,” *Technology Reports of Kansai University*, vol. 62, no. 06, pp. 2757–2767, Jul. 2020.
- [3] Z. A. Gumilang, “Implementasi naïve bayes classifier dan asosiasi untuk analisis sentimen data ulasan aplikasi e-commerce shopee pada situs google play,” Universitas Islam Indonesia, Yogyakarta, 2018.
- [4] F. Afshoh, “Analisa sentimen menggunakan naïve bayes untuk melihat persepsi masyarakat terhadap kenaikan harga jual rokok pada media sosial twitter,” Universitas Muhammadiyah Surakarta, Surakarta, 2017.
- [5] D. Alita and A. Rahman, “Pendeteksian Sarkasme pada Proses Analisis Sentimen Menggunakan Random Forest Classifier,” *Jurnal Komputasi*, vol. 8, no. 2, pp. 50–58, 2020.
- [6] S. L. Ting, W. H. Ip, and Albert. H. C. Tsang, “Is naïve bayes a good classifier for document classification?,” *International Journal of Software Engineering and Its Applications*, vol. 5, no. 3, pp. 37–46, Jul. 2011.
- [7] N. A. Verdikha, T. B. Adji, and A. E. Permanasari, “Komparasi metode oversampling untuk klasifikasi teks ujaran kebencian,” *Seminar Nasional Teknologi Informasi dan Multimedia*, vol. 6, no. 1, pp. 1.2-85-1,2-90, 2018.
- [8] G. James, D. Witten, T. Hastie, and R. Tibshirani, *An Introduction to Statistical Learning*, Second edition. New York: Springer, 2013.
- [9] S. Fanissa, M. A. Fauzi, and S. Adinugroho, “Analisis sentimen pariwisata di kota malang menggunakan metode naïve bayes dan seleksi fitur query expansion ranking,” *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, vol. 2, no. 8, pp. 2766–2770, Aug. 2018.
- [10] N. S. Wardani, A. Prahutama, and P. Kartikasari, “Analisis sentimen pemindahan ibu kota negara dengan klasifikasi naïve bayes untuk model bernoulli dan multinomial,” *Jurnal Gaussian*, vol. 9, no. 3, pp. 237–246, 2020.
- [11] N. Sulistiyowati and M. Jajuli, “Integrasi naïve bayes dengan teknik sampling SMOTE untuk menangani data tidak seimbang,” *Nuansa Informatika*, vol. 14, no. 1, pp. 34–37, 2020.
- [12] A. R. Purnajaya, W. A. Kusuma, and M. K. D. Hardhienata, “Prediksi interaksi pada jejaring bipartite senyawa dan protein pada data yang tidak seimbang,” Institut Pertanian Bogor, Bogor, 2019.