

ASPECT-BASED SENTIMENT ANALYSIS TERHADAP ULASAN APLIKASI FLIP MENGGUNAKAN PEMBOBOTAN TERM FREQUENCY-INVERSE DOCUMENT FREQUENCY (TF-IDF) DENGAN METODE KLASIFIKASI K-NEAREST NEIGHBORS (K-NN)

Ferda Ayu Dwi Putri Febrianti ^{1*}, Faqih Hamami ², Riska Yanu Fa'rifah ³

^{1*,2,3} Program Studi Sistem Informasi, Universitas Telkom, Kota Bandung, Provinsi Jawa Barat, Indonesia.

Email: ferdayuudwi@student.telkomuniversity.ac.id ^{1*}, faqihhamami@telkomuniversity.ac.id ², riskayanu@telkomuniversity.ac.id ³

Histori Artikel:

Dikirim 10 Agustus 2023; *Diterima dalam bentuk revisi* 20 Agustus 2023; *Diterima* 2 September 2023; *Diterbitkan* 10 September 2023. Semua hak dilindungi oleh Lembaga Penelitian dan Pengabdian Masyarakat (LPPM) STMIK Indonesia Banda Aceh.

Abstrak

Perkembangan transaksi online yang pesat di Indonesia telah meningkatkan kebutuhan akan solusi yang efisien dalam transfer antar bank. Namun, biaya yang terkait dengan transaksi semacam itu telah menjadi hambatan yang signifikan. Flip, sebuah perusahaan dengan visi untuk menjadi yang terdepan dalam layanan yang berfokus pada kepuasan pelanggan di seluruh dunia, menghadirkan solusi atas tantangan ini. Penelitian ini mengusulkan pendekatan Aspect-Based Sentiment Analysis dengan menggunakan algoritma K-Nearest Neighbors (K-NN) untuk menganalisis sentimen pengguna terhadap aspek-aspek kunci, yaitu kecepatan, keamanan, dan biaya dalam penggunaan aplikasi Flip. Hasil penelitian ini memberikan wawasan yang berharga yang dapat digunakan sebagai dasar untuk memberikan informasi, saran, dan rekomendasi kepada perusahaan, sehingga mereka dapat menciptakan solusi yang lebih baik dan meningkatkan pengalaman pengguna yang optimal. Hasil penelitian menunjukkan bahwa model K-NN mampu memprediksi sentimen pengguna dengan baik pada setiap aspek, dengan tingkat akurasi yang signifikan, yaitu kecepatan (73.04%), keamanan (86.05%), dan biaya (80.11%). Selain itu, penelitian ini juga membandingkan dua metode validasi model, yaitu metode split data sederhana dan K-fold cross validation. Meskipun metode split data sederhana memiliki akurasi rata-rata yang lebih tinggi, K-fold cross validation dianggap lebih unggul karena memberikan estimasi yang lebih akurat dan reliabel tentang kinerja model secara keseluruhan. Hasil analisis sentimen menunjukkan bahwa pengguna aplikasi Flip cenderung memberikan tanggapan negatif terhadap kecepatan dan keamanan, sementara mereka memberikan tanggapan positif terhadap biaya. Oleh karena itu, rekomendasi utama adalah bagi perusahaan PT Fliptech Lentera Inspirasi Pertiwi untuk meningkatkan aspek kecepatan dan keamanan guna meningkatkan kepuasan pengguna aplikasi Flip. Dengan demikian, layanan yang berfokus pada pelanggan ini akan terus memprioritaskan kepuasan pengguna sebagai tujuan utama.

Kata Kunci: ABSA; Ulasan Aplikasi; Kecepatan; Keamanan; Biaya; K-NN.

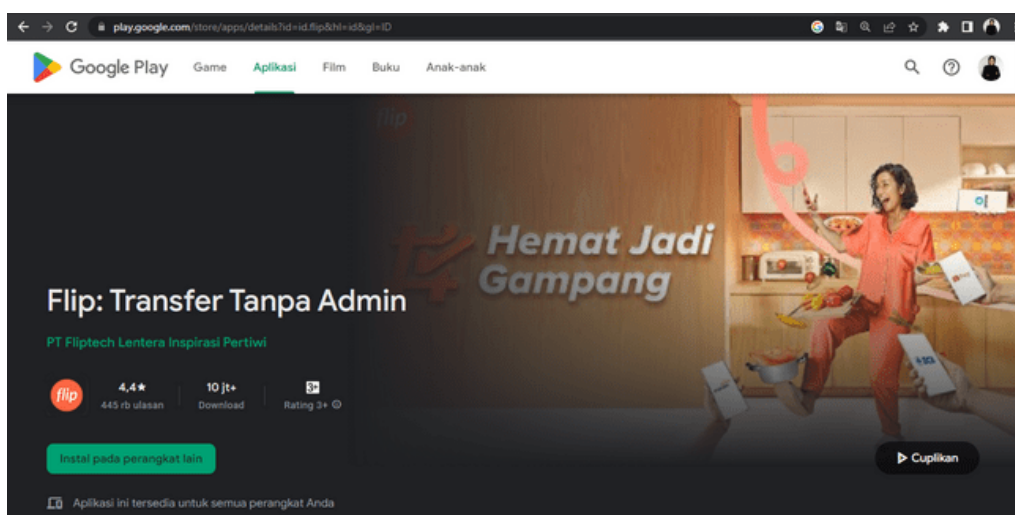
Abstract

The rapid growth of online transactions in Indonesia has increased the demand for efficient interbank transfer solutions. However, the costs associated with such transactions have become a significant obstacle. Flip, a company with a vision to become a global leader in customer satisfaction-driven services, offers a solution to this challenge. This study proposes an aspect-based sentiment analysis method using the K-Nearest Neighbors (K-NN) algorithm to analyze user sentiment on key aspects, namely speed, security, and the cost of using the Flip application. The results of this research provide valuable information that can be used as a basis to provide insights, suggestions and recommendations to businesses, so they can create better solutions and promote optimal user experience. The research results show that the K-NN model has the ability to predict user psychology well in all aspects, with a significant level of accuracy, specifically speed (73.04%), security (86, 05%) and costs (80.11%). In addition, this study also compares two model validation methods: simple data splitting method and K-Fold cross-validation. Although the simple data splitting method has a higher average accuracy, K-fold cross-validation is considered superior as it provides a more accurate and reliable estimate of the overall performance of the model. Sentiment analysis results show that Flip app users tend to give negative feedback on speed and security, while they give positive feedback on cost. Therefore, the main recommendation is that the company PT Fliptech Lentera Inspirasi Pertiwi improves the speed and security aspects to increase user satisfaction with the Flip application. Therefore, this customer-centric service will continue to prioritize user satisfaction as its primary goal.

Keyword: ABSA; App Reviews; Speed; Security; Cost; K-NN.

1. Pendahuluan

Saat ini transaksi *online* di Indonesia semakin meningkat, sehingga menyebabkan banyak nasabah memerlukan layanan transfer yang dapat diakses kapan saja. Untuk alasan ini, nasabah cenderung memilih *platform* atau bank yang menawarkan layanan transfer yang fleksibel agar dapat melakukan transaksi tanpa hambatan. Namun, transfer antar bank menjadi salah satu kendala, karena biaya yang dikenakan untuk setiap transaksi dapat mencapai Rp6.500. Bagi nasabah yang sering melakukan transfer, biaya itu bisa sangat membebani. Namun, Flip hadir sebagai solusi atas permasalahan ini. Flip adalah perusahaan teknologi keuangan yang beroperasi di Indonesia, menawarkan layanan transfer beda bank tanpa biaya, *top up e-wallet* tanpa potongan, dan transfer uang internasional dengan biaya yang lebih rendah. Hingga saat ini, Flip telah melayani lebih dari 12 juta pengguna serta 800 perusahaan dan usaha kecil menengah (UKM). Perusahaan ini juga menyediakan solusi manajemen keuangan bisnis terbaik bagi perusahaan dan pelaku bisnis [1]. Visi Flip adalah “Menjadi perusahaan dengan layanan paling *customer-centric* di dunia dan memungkinkan pengguna untuk melakukan transaksi keuangan yang adil dari mana saja kepada siapa pun” [2]. Perusahaan dengan layanan *customer-centric* menempatkan kepuasan pengguna sebagai prioritas utama dalam semua keputusan dan tindakan yang diambil. Dengan memprioritaskan kepuasan pengguna, perusahaan dapat membangun loyalitas pengguna yang kuat, meningkatkan retensi pengguna, dan memperoleh keunggulan kompetitif di pasar [3].



Gambar 1. Website Aplikasi Flip pada Google Play Store

Berdasarkan data pada Google Play Store seperti pada Gambar 1, hingga awal Maret 2023, khususnya di negara Indonesia, aplikasi Flip telah digunakan oleh lebih dari 10 juta pengguna, mendapatkan *rating* 4.4/5.0 serta ulasan berjumlah 445 ribu pada *platform* Google Play Store. Data tersebut membuktikan bahwa sebagai *start-up* lokal dalam 8 tahun terakhir Flip cukup diminati. Akan tetapi, suatu aplikasi tidak luput dari adanya kekurangan dan kelebihan masing-masing. Dimana hal tersebut dapat memunculkan berbagai respon atau tanggapan yang berbeda dari pengguna aplikasi seperti kepuasan dan kekecewaan terhadap aplikasi tersebut [4]. Google Play Store memiliki suatu fitur yang sangat berguna yakni kolom ulasan komentar terhadap suatu aplikasi yang menjadi salah satu tempat bagi pengguna untuk mengutarakan kepuasan, kekecewaan maupun opini terhadap aplikasi tersebut [5].

Analisis sentimen atau dikenal sebagai *opinion mining* merupakan metode yang digunakan secara otomatis dalam menemukan opini seseorang tentang suatu produk pada suatu teks. Metode ini menggunakan pendekatan *natural language processing*. Ada tiga cakupan berbeda untuk menerapkan analisis sentimen, yakni level dokumen, level kalimat, dan level aspek. Namun, analisis sentimen pada

level kalimat dan dokumen tidak cocok digunakan ketika suatu dokumen mengandung beberapa aspek yang berbeda. Contohnya ketika seseorang memberikan ulasan tentang salah satu jasa *tour travel*, ia mungkin memuji kualitas pelayanan yang diberikan oleh pihak *tour travel*, tetapi dalam ulasan yang sama ia juga mengkritik harga paket wisata *tour travel* yang sangat mahal. Maka dari itu, diperlukan pendekatan analisis sentimen berbasis aspek pada ulasan yang menyampaikan sentimen berbeda terhadap aspek yang juga berbeda dari suatu produk [6].

Aspect-Based Sentiment Analysis atau ABSA adalah bagian dari *opinion mining* yang memungkinkan pengguna memperoleh informasi lebih rinci tentang aspek yang dibahas dalam ulasan. Pendekatan ABSA terdiri dari dua tahapan, yaitu *filtering statements* dan *extracting sentiments* [7]. Dengan ABSA, dapat mengidentifikasi aspek tertentu yang dibahas dalam data dan menentukan sentimen terkait dengan masing-masing aspek [8]. Dengan memahami respons atau tanggapan yang diberikan oleh pengguna terhadap beberapa aspek dan mengetahui sentimen terkait masing-masing aspek, dapat membantu dalam memperbaiki dan meningkatkan produk, layanan, aplikasi, atau kinerja yang sesuai dengan umpan balik pengguna terhadap aspek tersebut [9]. Sehingga dari penelitian ini diharapkan dapat diketahui sentimen pengguna berbasis aspek pada objek aplikasi Flip terhadap kualitas layanan aplikasi Flip dari *review* atau komentar berupa opini yang diekspresikan dalam teks yang diberikan pengguna melalui *platform* Google Play Store dan mencari tahu aspek yakni kecepatan, keamanan dan biaya yang mendapat penilaian positif, negatif, ataupun tidak ada dari pengguna agar Flip dapat melakukan evaluasi untuk meningkatkan layanan aplikasi sesuai yang diutarakan oleh pengguna. *Aspect-Based Sentiment Analysis* dapat diaplikasikan menggunakan metode klasifikasi untuk mempermudah dalam pengelompokan data yang dikategorikan dari beberapa aspek dengan metode K-NN (*K-Nearest Neighbor*) [10]. Metode K-NN digunakan pada proses klasifikasi karena memiliki proses dan logika algoritma yang sederhana, mudah dipahami dan diimplementasikan, serta memiliki nilai akurasi yang lumayan tinggi untuk memecahkan tantangan dalam klasifikasi [11]. Selain itu, pemilihan algoritma K-NN memiliki justifikasi yang kuat berdasarkan karakteristik khas dari algoritma ini. Beberapa alasan mengapa algoritma K-NN merupakan pilihan yang tepat adalah sebagai berikut:

- 1) Algoritma K-NN menerapkan prinsip "*Bird of a Feather*" dalam menentukan lokasi yang tepat bagi data baru. Algoritma K-NN mengasumsikan bahwa objek-objek serupa cenderung berada dalam jarak yang dekat atau bertetangga. Dengan kata lain, data yang memiliki kesamaan cenderung berada dekat satu sama lain.
- 2) K-NN memanfaatkan seluruh *dataset* yang ada dan mengategorikan data baru berdasarkan tingkat kesamaan atau fungsi jarak. Data baru selanjutnya ditempatkan ke dalam kelas di mana sebagian besar tetangga data berada.
- 3) K-NN memiliki sedikit *hyperparameter*. Hanya diperlukan nilai *k* dan metrik jarak, yang jumlahnya lebih sedikit jika dibandingkan dengan algoritma *machine learning* lainnya.
- 4) Algoritma K-NN termasuk kelompok algoritma non-parametrik, tidak memerlukan asumsi tentang distribusi data. Dalam analisis sentimen berbasis aspek, karakteristik sentimen terhadap aspek-aspek tertentu mungkin memiliki distribusi yang kompleks. Oleh karena itu, pendekatan non-parametrik dapat lebih adaptif dalam menangani variasi yang kompleks.
- 5) Dalam analisis sentimen berbasis aspek, frekuensi sentimen positif, negatif, dan netral terhadap suatu aspek mungkin tidak seimbang. Algoritma K-NN dapat mengatasi masalah ketidakseimbangan tersebut dengan mempertimbangkan kelas-kelas yang lebih sedikit secara adil dalam proses pemilihan tetangga terdekat.
- 6) K-NN mengambil keputusan berdasarkan mayoritas label kelas dari tetangga terdekat. Dalam analisis sentimen berbasis aspek, sentimen terkait suatu aspek cenderung berkumpul bersama, sehingga penting untuk mempertimbangkan sentimen mayoritas terdekat dari aspek yang sama. Maka dari itu, K-NN memiliki keterbacaan dan interpretabilitas yang baik.
- 7) K-NN memiliki ketahanan terhadap data pencilan (*outliers*) karena tidak membangun model prediksi yang kompleks. Dalam analisis sentimen berbasis aspek, adanya komentar atau ulasan yang tidak representatif terhadap suatu aspek dapat berdampak pada hasil sentimen. K-NN dapat menangani data-pencililan dengan tidak terlalu dipengaruhi oleh titik-titik yang jauh.

- 8) Algoritma K-NN bersifat *lazy learning*, untuk membuat model tidak memakai titik *data training*. Dalam K-NN tidak ada tahap *training*, sehingga seluruh *data training* diaplikasikan pada tahap *testing* menyebabkan proses *training* lebih cepat tetapi mengakibatkan tahap *testing* lebih lambat, karena K-NN memerlukan waktu lebih lama untuk menganalisis setiap titik data. Terlebih lagi, proses ini memerlukan penggunaan memori yang lebih besar untuk menyimpan *data training*.

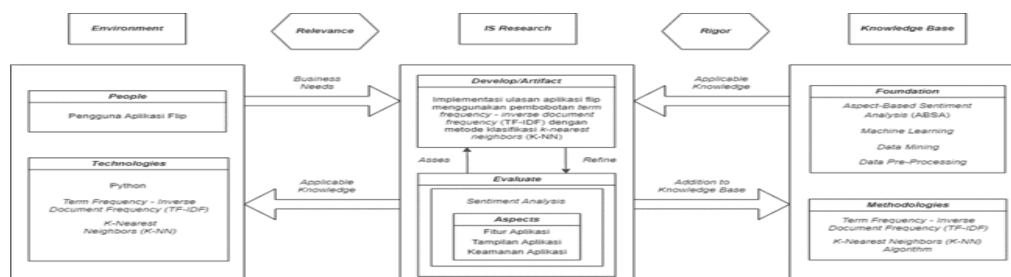
Penggunaan metode *Term Frequency-Inverse Document Frequency* dikarenakan memiliki kesederhanaan dimana prosesnya berdasarkan pendekatan pembobotan yang sederhana, juga kemudahan dalam implementasi, adaptasi dan proses *learning*. Pembobotan kata adalah proses pemberian bobot untuk setiap kata yang terdapat dalam suatu dokumen. Dalam metode TF-IDF, *Term Frequency* lebih berfokus pada istilah yang sering muncul dalam suatu dokumen sedangkan *Inverse Document Frequency* lebih berfokus pada pemberian bobot rendah untuk istilah yang muncul dalam banyak dokumen [12]. Pemilihan metode dan algoritma tersebut didasarkan pada penelitian terdahulu, dimana K-NN merupakan salah satu metode yang dipakai untuk melakukan analisis sentimen terhadap produk atau aplikasi. Pada penelitian tersebut didapatkan tingkat akurasi klasifikasi analisis sentimen pada aplikasi Flip mencapai 76.68% menunjukkan keefektifan metode ini. Dalam evaluasi klasifikasi, nilai presisi untuk kelas data dengan ulasan positif mencapai 82.67% artinya rasio data yang benar mengandung ulasan positif relatif tinggi dibandingkan dengan keseluruhan data yang diprediksi mengandung ulasan positif. Nilai *Recall* mencapai 86.92%, menunjukkan bahwa rasio data yang diprediksi mengandung ulasan positif dibandingkan dengan keseluruhan data yang benar mengandung ulasan positif cukup tinggi [5].

Penelitian ini mengimplementasikan *Aspect-Based Sentiment Analysis* pada ulasan aplikasi Flip di Google Play Store menggunakan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan metode klasifikasi *K-Nearest Neighbors* (K-NN). Beberapa hal yang dijadikan sebagai batasan pada penelitian ini diantaranya yaitu menggunakan data ulasan aplikasi Flip yang diambil dari platform Google Play Store pada tahun 2023 dimana hanya memakai data ulasan berbahasa Indonesia, menggunakan tiga polaritas sentimen yakni tidak ada aspek, sentimen positif dan sentimen negatif, *Aspect-Based Sentiment Analysis* yang dilakukan tertuju pada ulasan aplikasi Flip dilihat dari aspek kecepatan, keamanan dan biaya. Alasan dipilihnya ketiga aspek tersebut adalah karena terdapat banyak faktor yang memengaruhi masyarakat untuk beralih dari transaksi konvensional ke digital. Beberapa di antaranya termasuk kemudahan penggunaan, keamanan, kecepatan, kredibilitas, faktor kepercayaan, dan manfaat biaya dari produk *digital wallet* tersebut [13]. Tujuan penelitian ini ialah menganalisis *Aspect-Based Sentiment Analysis* Terhadap Ulasan Aplikasi Flip Menggunakan Pembobotan *Term Frequency - Inverse Document Frequency* (TF-IDF) Dengan Metode Klasifikasi *K-Nearest Neighbors* (K-NN).

2. Metode Penelitian

2.1 Model Konseptual

Model konseptual merupakan rancangan yang menunjukkan hubungan logis antara faktor atau variabel yang telah diidentifikasi dengan melibatkan beberapa aspek dan proses aktivitas. Berikut model konseptual dalam penelitian:

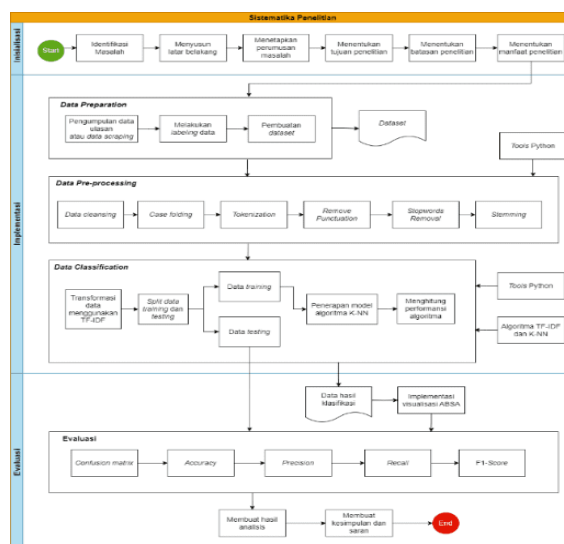


Gambar 2. Model Konseptual

Model konseptual terbagi menjadi tiga bagian diantaranya *Environment* (lingkungan penelitian), *IS Research* (penelitian pada bidang Sistem Informasi), dan *Knowledge Base* (basis pengetahuan atau dasar ilmu yang digunakan). *Environment* terbagi menjadi dua bagian yaitu *People* (orang atau aktor yang ikut terlibat dalam penelitian) dan *Technologies* (teknologi yang digunakan dalam penelitian). *People* mencakup pengguna aplikasi Flip, sedangkan *Technologies* yang digunakan Python, pembobotan TF-IDF dan algoritma K-NN. Penelitian ini akan menghasilkan penerapan implementasi TF-IDF dan K-NN untuk *sentiment analysis* dengan *aspects* yakni kecepatan, keamanan dan biaya pada ulasan aplikasi Flip di Google Play Store, dengan menerapkan metode *Aspect-Based Sentiment Analysis* (ABSA), *Machine Learning*, *Data Mining*, *Data Pre-Processing*. Tujuan akhir penelitian adalah menyelesaikan masalah dalam mengklasifikasikan ulasan pengguna aplikasi Flip menjadi tiga kelas yakni kelas tidak ada aspek, sentimen positif dan sentimen negatif, sehingga menambah *insight* baru bagi perusahaan untuk mengembangkan layanan aplikasi sesuai kebutuhan pengguna pada aspek terkait.

2.2 Sistematika Penelitian

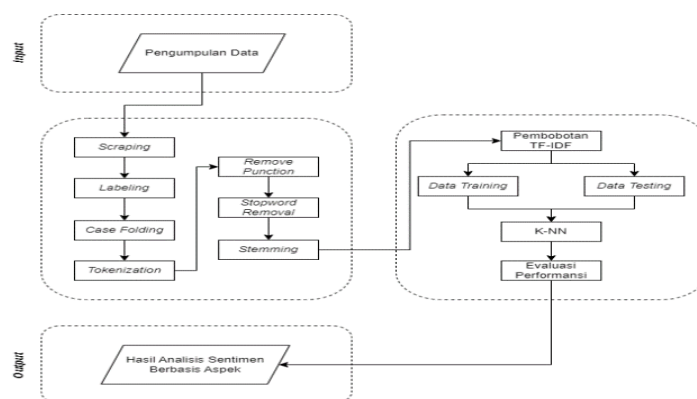
Sistematika penelitian bertujuan mengetahui proses yang Akan dilakukan selama penelitian agar mencapai solusi dari setiap permasalahan yang ada. Sistematika penelitian dibuat dalam diagram alur (*flowchart*) yang terbagi menjadi tahap inisialisasi, implementasi, dan evaluasi.



Gambar 3. Sistematika Penelitian

2.3 Tahapan Analisis

Tahapan analisis dimulai dari data *scraping*, *labeling* data, kemudian tahap *pre-processing* dengan melakukan *cleansing* data, *case folding*, *tokenization*, *remove punctuation*, *stopword removal*, dan *stemming*. Kemudian data dibagi menjadi *training* dan *testing*. Selanjutnya dilakukan pembobotan TF-IDF lalu klasifikasi menggunakan algoritma K-NN. Terakhir, dilakukan perhitungan evaluasi dengan metode *confusion matrix*.



Gambar 4. Tahapan Analisis

1) *Labeling Data*

Data yang digunakan adalah data ulasan aplikasi Flip di Indonesia pada *platform* Google Play Store. Kumpulan data disimpan ke dalam *file* dengan format csv. Total data mentah yang terkumpul sebanyak 121.834 ulasan. *Dataset* kemudian diberi label yang terdiri dari tidak ada aspek, positif, dan negatif. Data yang telah diberikan label berjumlah 13.500 data. Pelabelan dilakukan secara manual untuk seluruh data sebelum dilakukannya data *pre-processing*. Data yang dipakai setidaknya memiliki minimal 1 aspek di dalamnya. Berikut penjelasan ketiga aspek:

- Aspek kecepatan merujuk pada waktu pemrosesan aplikasi dalam memberikan respons kepada pengguna serta seberapa cepat aplikasi dapat menangani berbagai operasi keuangan seperti transaksi finansial.
- Aspek keamanan merujuk pada keamanan informasi sensitif pengguna, kerahasiaan data saat proses transaksi finansial, keamanan pada setiap fitur seperti *login*, *register*, verifikasi akun, juga *refund* maupun keamanan aplikasi pada saat digunakan.
- Aspek biaya merujuk pada berbagai jenis biaya yang dikeluarkan dalam proses transaksi finansial, khususnya yang berkaitan dengan administrasi dan biaya dalam penggunaan aplikasi.

Ketiga aspek tersebut diberi label yang dikategorikan dalam tiga kelas yaitu 0 (tidak ada aspek), 1 (positif) dan 2 (negatif). Berikut penjelasan ketiga label kelas:

- Label 0 kelas tidak ada aspek. Ulasan tidak mengekspresikan pandangan positif maupun negatif secara jelas berdasarkan aspek tertentu.
- Label 1 kelas sentimen positif. Adanya pandangan baik, pengalaman memuaskan, keunggulan, atau manfaat yang dirasakan pengguna terhadap aplikasi. Ulasan berisi ungkapan kepuasan, pujian, atau apresiasi terhadap penggunaan aplikasi.
- Label 2 kelas sentimen negatif. Adanya pandangan buruk, pengalaman tidak memuaskan, kelemahan, atau masalah yang dirasakan pengguna terhadap aplikasi. Ulasan mengandung kritik, ketidakpuasan, atau ketidakcocokan terhadap penggunaan aplikasi.

2) *Pre-processing Data*

Setelah pelabelan data, dilakukan pengolahan data dengan proses data *pre-processing* yang terdiri dari beberapa tahapan yaitu data *cleansing*, *case folding*, *tokenization*, *remove punctuation*, *stopwords removal*, dan *stemming*.

3) Penentuan Nilai K

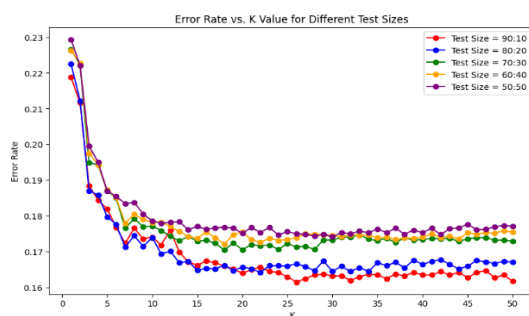
Dalam penelitian, metode yang digunakan adalah K-NN untuk melakukan klasifikasi. Pada tahap klasifikasi, data dipisah menjadi dua proses, yaitu pelatihan dan pengujian. Proses pelatihan untuk menghasilkan model, lalu proses pengujian menggunakan perhitungan jarak *cosine similarity* sebagai ukuran kesamaan antara data yang akan diklasifikasikan dengan data pelatihan. Data dibagi menjadi rasio pembagian 50:50, 60:30, 70:30, 80:20, dan 90:10. Rasio tersebut diuji untuk mengukur

error rate dan akurasi. Rasio pembagian dengan *error rate* terendah serta tingkat akurasi tertinggi akan digunakan dalam penelitian ini.

Tabel 1. Error Rate Terendah Pada Kelima Rasio

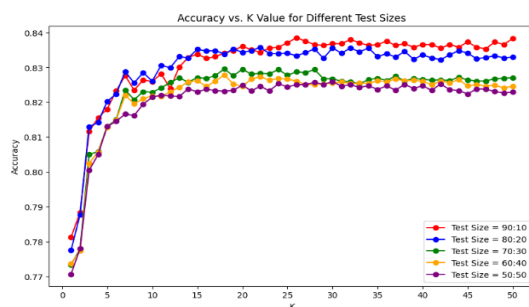
Rasio	K Terbaik	Error Rate Terendah
50:50	30	17.41%
60:40	18	17.21%
70:30	18	17.04%
80:20	22	16.41%
90:10	26	16.14%

Rasio 90:10 memiliki nilai *error rate* paling rendah (16.14%) dibandingkan dengan rasio lainnya dengan K=26.



Gambar 5. Perbandingan Nilai Error Rate

Nilai K yang dipilih adalah 26 dilihat dari nilai *error rate* terendah pada rentang nilai K yaitu 25 sampai 50. Untuk memvalidasi kesimpulan tersebut, dilakukan evaluasi terhadap akurasi model menggunakan data pelatihan dengan menggunakan nilai K=26.



Gambar 6. Perbandingan Nilai Akurasi

Pada rasio 90:10 jika nilai $K \geq 25$, akan diperoleh tingkat akurasi lebih besar dari 83.00%. Oleh karena itu, dapat disimpulkan hasil penentuan nilai K=26 benar-benar menghasilkan tingkat akurasi yang tinggi (83.85%). Akurasi mencakup semua aspek yang ada dalam label y kemudian diambil rata-rata dari akurasi-akurasi tersebut.

Tabel 2. Akurasi Tertinggi Pada Kelima Rasio

Rasio	K Terbaik	Akurasi Tertinggi
50:50	30	82.58%
60:40	18	82.78%
70:30	18	82.95%
80:20	22	83.58%
90:10	26	83.85%

4) TF-IDF

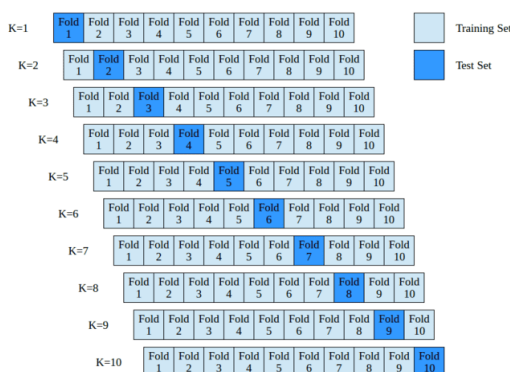
Setelah tahap data *pre-processing*, selanjutnya memberikan bobot pada setiap kata dalam data ulasan aplikasi Flip. Proses pemberian bobot dilakukan dengan metode TF-IDF. Nilai perhitungan bobot yang diperoleh dari setiap kata digunakan sebagai *input* dalam proses klasifikasi.

5) GridSearch

Pengujian dengan GridSearch bertujuan mencari kombinasi parameter model yang menghasilkan performa optimal [14]. Parameter *weight*, *metric*, dan *algorithm* pada model K-NN dites secara sistematis untuk menemukan nilai terbaik. Pengujian dilakukan pada rasio 90:10 serta dengan K=26 untuk mengurangi total waktu yang dibutuhkan dalam proses GridSearch saat melatih dan menguji model pada setiap iterasi. Sehingga mengefisienkan penggunaan waktu komputasi dan sumber daya yang tersedia. Hasil GridSearch menjadi acuan dalam penggunaan parameter model pada metode *split* data sederhana dan *K-fold cross validation*.

6) K-Fold Cross Validation

Pengujian dengan *K-fold cross validation* untuk mendapatkan estimasi dari model yang dibuat dengan perulangan proses sebanyak *n_splits*, di mana nilai *n_splits* adalah 10 *folds* [15]. Data dibagi menjadi 10 *subset* yang sama, setiap *fold* akan menjadi data validasi satu kali. Memakai K=10 yang berarti data pelatihan akan mencakup 90% dari total data, sedangkan data pengujian akan mencakup 10% dari total data. Pembagian secara acak untuk menghindari bias dalam pengujian. Pada setiap iterasi, salah satu *fold* akan digunakan sebagai data pengujian (*test*), sedangkan 9 *folds* lainnya akan digunakan sebagai data pelatihan (*train*). Proses diulang sebanyak K kali dengan ukuran yang sama, sehingga setiap *fold* pernah menjadi data pengujian.



Gambar 7. K-Fold Cross Validation

7) Evaluasi Performansi

Evaluasi performansi untuk mengetahui kinerja proses yang dilakukan dengan *confusion matrix* sebagai metode terakhir yang bertujuan menguji *accuracy*, *precision*, *recall*, dan *f1-score* dari model klasifikasi. *Confusion matrix* memberi gambaran tentang seberapa baik kinerja model hingga dapat mengklasifikasikan data ke dalam kelas yang benar. Evaluasi *confusion matrix* memberi informasi tentang seberapa banyak data uji yang terklasifikasi dengan benar dan salah. Selain itu, evaluasi juga menganalisis hasil prediksi model terhadap sentimen yang ada [16].

3. Hasil dan Pembahasan

3.1 Grid Search

Skenario pengujian dengan *GridSearch* untuk menemukan kombinasi parameter model yang memberi hasil paling optimal sehingga diketahui akurasi tertinggi pada setiap aspek. Memakai rasio 90:10 dengan $K = 26$. Parameter yang diuji yakni *estimator algorithm*, *metric*, dan *weights*.

Estimator algorithm	Estimator metric	Estimator weights	Akurasi			Rata-rata
			Kecapatan	Kemampuan	Riwayat	
auto	euclidean	uniform	78.22%	85.55%	86.00%	83.25%
		distance	78.37%	85.62%	85.62%	83.20%
	Rata-Rata		78.29%	85.59%	85.81%	83.23%
	cosine	uniform	78.00%	86.59%	86.96%	83.85%
		distance	78.44%	86.66%	86.81%	83.97%
	Rata-rata		78.22%	86.62%	86.88%	83.91%
	minkowski	uniform	78.22%	85.55%	86.00%	83.25%
		distance	78.37%	85.62%	85.62%	83.20%
	Rata-Rata		78.29%	85.59%	85.81%	83.23%
kd_tree	euclidean	uniform	78.22%	85.55%	86.00%	83.25%
		distance	78.37%	85.62%	85.62%	83.20%
	Rata-rata		78.29%	85.59%	85.81%	83.23%
	cosine	uniform	78.00%	86.59%	86.96%	83.85%
		distance	78.44%	86.66%	86.81%	83.97%
	Rata-rata		78.22%	86.62%	86.88%	83.91%
	minkowski	uniform	78.22%	85.55%	86.00%	83.25%
		distance	78.37%	85.62%	85.62%	83.20%
	Rata-rata		78.29%	85.59%	85.81%	83.23%
ball_tree	euclidean	uniform	78.22%	85.55%	86.00%	83.25%
		distance	78.37%	85.62%	85.62%	83.20%
	Rata-rata		78.29%	85.59%	85.81%	83.23%
	cosine	uniform	78.00%	86.59%	86.96%	83.85%
		distance	78.44%	86.66%	86.81%	83.97%
	Rata-rata		78.22%	86.62%	86.88%	83.91%
	minkowski	uniform	78.22%	85.55%	86.00%	83.25%
		distance	78.37%	85.62%	85.62%	83.20%
	Rata-rata		78.29%	85.59%	85.81%	83.23%

Gambar 8. Akurasi Pada Gridsearch

Pengujian *estimator metric cosine* dengan kombinasi *parameter estimator weights distance* memiliki rata-rata akurasi keseluruhan tertinggi (83.91%). Pengujian dengan *estimator metric euclidean* dan *minkowski* memiliki rata-rata akurasi keseluruhan yang sama (83.23%). Dengan demikian, jika dilihat dari *estimator metric*, pengujian dengan *estimator metric cosine* dikatakan paling baik dalam hal akurasi secara keseluruhan. *Estimator algorithm (auto, kd_tree, ball_tree)* menghasilkan hasil akurasi yang sama persis, menunjukkan pemakaian *estimator algorithm* yang berbeda pada model tidak memiliki pengaruh yang signifikan pada akurasi model.

3.2 Split Data Sederhana

Skenario pengujian dengan *split* data sederhana untuk mengetahui akurasi yang didapatkan ketika menggunakan *split* data *train-test* dengan *library scikit-learn*. Memakai $K=26$, rasio 90:10, *metric=cosine*, serta parameter model yang serupa dengan hasil pada pengujian *GridSearch*.

Split Data	Akurasi			Rata-rata
	Kecepatan	Kemanan	Biaya	
Train set	99.94%	99.99%	99.95%	99.96%
Test set	78.44%	86.66%	86.81%	83.97%
Rata-rata	89.19%	93.32%	93.38%	91.96%

Gambar 9. Akurasi Pada *Split* Data Sederhana

Rata-rata akurasi *test set* dari ketiga aspek (83.97%). Model mengalami penurunan kinerja saat dihadapkan pada data baru yang tidak pernah dilihat sebelumnya. Penurunan ini disebabkan oleh *overfitting* atau karakteristik yang berbeda antara data *training* dan data *testing*. Rata-rata akurasi keseluruhan (91.96%). Model memiliki tingkat keberhasilan yang baik secara keseluruhan dalam memprediksi label kecepatan, keamanan, dan biaya. Model memiliki kinerja yang baik dalam memprediksi aspek kecepatan, keamanan, dan biaya, terutama pada aspek keamanan dan biaya. Namun, terdapat perbedaan yang signifikan antara akurasi *train set* dan *test set*, yang mengindikasikan adanya *overfitting*.

3.3 K-Fold Cross Validation

Pada skenario pengujian dengan *K-fold cross validation*, data dibagi menjadi *K subset* atau lipatan (*folds*) yang sama ukurannya. *Dataset* dibagi menjadi 10 *subset* ($K=10$). *Subset* pertama digunakan sebagai *data training*, sisanya digunakan sebagai *data testing*. Setelah itu, dihitung rata-rata akurasi dari 10 iterasi tersebut, sehingga memberikan perkiraan kinerja model secara keseluruhan.

Split Data	Akurasi		
	Kecepatan	Kemanan	Biaya
Fold 1	60.81%	84.00%	71.77%
Fold 2	70.22%	86.22%	78.22%
Fold 3	55.18%	84.14%	78.96%
Fold 4	68.29%	85.11%	81.70%
Fold 5	58.88%	84.44%	70.07%
Fold 6	66.59%	85.18%	68.66%
Fold 7	61.11%	85.40%	73.18%
Fold 8	60.44%	84.74%	70.00%
Fold 9	59.85%	84.29%	70.59%
Fold 10	60.96%	84.74%	70.88%
Rata-rata	62.23%	84.82%	73.40%

Gambar 10. Akurasi Pada *K-Fold Cross Validation*

Model memiliki kinerja yang baik dalam memprediksi aspek keamanan, sedangkan kinerjanya cenderung rendah dalam memprediksi aspek kecepatan. Secara keseluruhan, model memiliki tingkat keberhasilan yang moderat dalam memprediksi label kecepatan, keamanan, dan biaya.

3.4 Evaluasi Performansi

Hasil evaluasi performansi menggunakan data tahap klasifikasi dengan $K=26$, *metric=cosine*, *algorithm=auto*, *weights=distance*, serta rasio 90:10 atau pembagian data sebanyak 90% data *test* dan 10% data *train*. Dengan *confusion matrix*, didapatkan hasil *Accuracy*, *Precision*, *Recall*, dan *F1-Score* pada aspek kecepatan, keamanan dan biaya. *Accuracy* mengukur sejauh mana model dapat mengklasifikasikan dengan benar keseluruhan data, *precision* mengukur sejauh mana model dapat mengidentifikasi dengan

benar setiap kelas, *Recall* mengukur sejauh mana model dapat mendeteksi dengan benar setiap kelas, dan *F1-Score* menggabungkan *precision* dan *recall* menjadi satu metrik mencerminkan keseimbangan antara keduanya.

Aspek Kecepatan								
Confusion Matrix		Predicted			Accuracy	Precision	Recall	F1-Score
		None	Positif	Negatif				
Actual	None	734	18	29	0.78	0.78	0.94	0.85
	Positif	102	132	11		0.74	0.54	0.62
	Negatif	102	29	193		0.83	0.60	0.69
Macro Average						0.78	0.69	0.72

Gambar 11. Evaluasi Performansi Aspek Kecepatan

Klasifikasi aspek kecepatan, kelas 0 (tidak ada aspek) memiliki tingkat akurasi yang tinggi, dengan 734 prediksi benar. Kelas 1 (positif) memiliki *false negatives* yang cukup tinggi, dengan 102 kasus yang salah diklasifikasikan sebagai kelas 0. Kelas 2 (negatif) memiliki *false positives* yang cukup tinggi, dengan 29 kasus yang salah diklasifikasikan sebagai kelas 0. Model memiliki tingkat ketepatan yang baik dalam mengklasifikasikan kelas 0, namun memiliki tingkat ketepatan yang lebih rendah dalam mengklasifikasikan kelas 1 dan kelas 2. Model juga memiliki tingkat deteksi yang baik untuk kelas 0, tetapi lebih rendah untuk kelas 1 dan kelas 2.

Aspek Keamanan								
Confusion Matrix		Predicted			Accuracy	Precision	Recall	F1-Score
		None	Positif	Negatif				
Actual	None	1074	2	24	0.87	0.88	0.98	0.92
	Positif	60	14	2		0.88	0.18	0.30
	Negatif	92	0	82		0.76	0.47	0.58
Macro Average						0.84	0.54	0.60

Gambar 12. Evaluasi Performansi Aspek Keamanan

Klasifikasi aspek keamanan, kelas 0 (tidak ada aspek) memiliki tingkat akurasi yang tinggi, dengan 1074 prediksi benar. Kelas 1 (positif) memiliki *false negatives* yang cukup tinggi, dengan 60 kasus yang salah diklasifikasikan sebagai kelas 0. Kelas 2 (negatif) memiliki *false positives* yang cukup tinggi, dengan 24 kasus yang salah diklasifikasikan sebagai kelas 0. Model memiliki tingkat ketepatan yang baik dalam mengklasifikasikan kelas 0, namun memiliki tingkat ketepatan yang lebih rendah dalam mengklasifikasikan kelas 1 dan kelas 2. Model juga memiliki tingkat deteksi yang baik untuk kelas 0, tetapi lebih rendah untuk kelas 1 dan kelas 2.

Aspek Biaya								
Confusion Matrix		Predicted			Accuracy	Precision	Recall	F1-Score
		None	Positif	Negatif				
Actual	None	459	106	1	0.87	0.91	0.81	0.86
	Positif	38	710	0		0.85	0.95	0.89
	Negatif	10	23	3		0.75	0.08	0.15
Macro Average						0.83	0.61	0.63

Gambar 13. Evaluasi Performansi Aspek Biaya

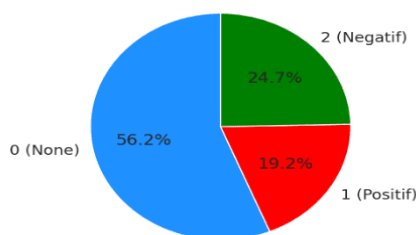
Klasifikasi aspek biaya, kelas 0 (tidak ada aspek) memiliki tingkat akurasi yang cukup tinggi, dengan 459 prediksi benar. Kelas 1 (positif) memiliki *false negatives* yang cukup tinggi, dengan 38 kasus yang salah diklasifikasikan sebagai kelas 0. Kelas 2 (negatif) memiliki *false negatives* dan *false positives* yang cukup signifikan. Model memiliki tingkat ketepatan yang baik dalam mengklasifikasikan kelas 0 dan kelas 1, tetapi memiliki tingkat ketepatan yang rendah dalam mengklasifikasikan kelas 2. Model juga memiliki tingkat deteksi yang baik untuk kelas 1, tetapi rendah untuk kelas 2.

3.5 Analisis Sentimen Berbasis Aspek

Analisis berbasis aspek divisualisasikan menggunakan grafik *pie chart* dan *word cloud*.

1) Aspek Kecepatan

Persentase Jumlah 0, 1, dan 2 pada aspek Kecepatan



Gambar 14. *Pie Chart* Aspek Kecepatan

Hasil kelas 0 (tidak ada aspek) sebesar 7.584 atau sekitar 56.2%, lalu kelas 1 (positif) sebesar 2.586 atau sekitar 19.2% dan kelas 2 (negatif) sebesar 3.330 atau sekitar 24.7%. Tidak ada aspek memiliki jumlah lebih besar dari pada kedua kelas lainnya. Namun kelas negatif cenderung lebih tinggi dari kelas positif. Dapat disimpulkan pengguna Flip lebih banyak menanggapi negatif terhadap kecepatan aplikasi Flip.

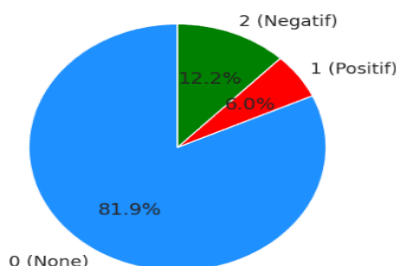


Gambar 15. *Wordcloud* Aspek Kecepatan

None – Tidak memiliki Aspek: Kata “Flip” dan “Transfer” merupakan kata yang sering muncul, menunjukkan topik utama yang dibicarakan adalah transfer uang melalui aplikasi Flip. Selain itu, kata “Biaya” tidak termasuk ke aspek kecepatan sehingga dikategorikan sebagai *None* atau tidak memiliki aspek kecepatan. **Positif**: Kata “Cepat” paling sering muncul, menunjukkan kecepatan adalah aspek yang sangat diunggulkan. Kemudian kata “Flip”, “Transfer”, “Proses”, dan “Transaksi” menunjukkan kecepatan dalam melakukan transfer uang dan proses transaksi adalah perhatian utama bagi pengguna aplikasi Flip dalam menggambarkan pengalaman yang efisien saat menggunakan aplikasi. **Negatif**: Kata “Transfer”, “Flip” dan “Proses” paling banyak muncul, namun terdapat penurunan jumlah kata “Cepat” dibandingkan dengan *word cloud* sebelumnya, menunjukkan adanya aspek kecepatan yang menjadi perhatian negatif. Selain itu, kata “Menit”, “Jam” dan “Nunggu” menunjukkan bahwa pengguna merasa transfer atau proses transaksi memakan waktu yang lama.

2) Aspek Keamanan

Persentase Jumlah 0, 1, dan 2 pada aspek Keamanan



Gambar 16 Pie Chart Aspek Keamanan

Hasil kelas 0 (tidak ada aspek) sebesar 11.052 atau sekitar 81.9%, lalu kelas 1 (positif) sebesar 806 atau sekitar 6.0% dan kelas 2 (negatif) sebesar 1.642 atau sekitar 12.2%. Tidak ada aspek memiliki jumlah lebih besar dari pada kedua kelas lainnya. Namun kelas negatif cenderung lebih tinggi dari kelas positif. Dapat disimpulkan pengguna Flip lebih banyak menanggapi negatif terhadap keamanan aplikasi Flip.

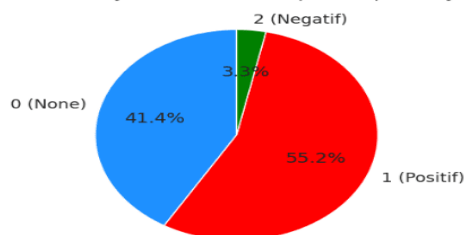


Gambar 17. Wordcloud Aspek Keamanan

None – Tidak memiliki Aspek: Kata “Flip” dan “Transfer” memiliki jumlah frekuensi yang tinggi. Kemudian terdapat kata seperti “Biaya”, “Hemat”, dan “Cepat” menunjukkan bahwa kata-kata tersebut tidak termasuk ke aspek keamanan sehingga dikategorikan sebagai *None* atau tidak memiliki aspek keamanan. Positif: Kata “Flip” dan “Transfer” tetap muncul dalam jumlah yang signifikan, menunjukkan bahwa pengguna mengaitkan aspek keamanan dengan transfer dan transaksi melalui aplikasi Flip. Terlihat kata “Aman” juga sangat sering disebutkan, menunjukkan pengguna merasa aplikasi Flip memberikan tingkat keamanan yang memadai dalam proses transfer dan transaksi. Negatif: Kata “Flip” dan “Transfer” adalah kata yang banyak disebutkan. Kata “Saldo”, “Rekening”, “Kirim”, dan “Gagal” mengindikasikan adanya masalah terkait keamanan saldo, keamanan rekening, keamanan dalam pengiriman dana, dan kegagalan transaksi. Kata “Kecewa” dan “Ribet” menunjukkan adanya kekecewaan pengguna terkait tingkat keamanan aplikasi Flip dikarenakan keamanan yang buruk.

3) Aspek Biaya

Persentase Jumlah 0, 1, dan 2 pada aspek Biaya



Gambar 18. Pie Chart Aspek Biaya

Jurnal Indonesia : Manajemen Informatika dan Komunikasi, Volume 4 No 3, September (2023), pp. 1858-1873
DOI : <https://doi.org/10.35870/jimik.v4i3.429>

sederhana. Meskipun *K-fold cross validation* menghasilkan akurasi yang lebih rendah pada setiap iterasi, metode ini lebih baik karena memberikan estimasi yang lebih akurat dan reliabel tentang kinerja model secara keseluruhan. Pengguna Flip lebih banyak menanggapi negatif terhadap kecepatan, negatif terhadap keamanan, serta positif terhadap biaya aplikasi Flip. Setiap aspek memiliki kecenderungan kata yang berbeda. Pada aspek kecepatan negatif kata “Menit”, “Jam” dan “Nunggu” menunjukkan pengguna merasa transfer atau proses transaksi memakan waktu yang lama. Pada aspek keamanan negatif kata “Saldo”, “Rekening”, “Kirim”, dan “Gagal” mengindikasikan adanya masalah terkait keamanan saldo, keamanan rekening, keamanan dalam pengiriman dana, dan kegagalan transaksi. Pada aspek biaya positif kata “Biaya”, “Hemat”, “Bank”, dan “Aplikasi” menunjukkan adanya perhatian pengguna terhadap biaya yang hemat dalam penggunaan aplikasi Flip sebagai layanan perbankan. Kata “Gratis”, “Proses”, “Kirim”, “Suka”, “Murah”, dan “Sukses” juga menunjukkan persepsi positif biaya sehingga adanya kepuasan pengguna.

5. Daftar Pustaka

- [1] Flip.id, “Tentang Flip,” 2023. <https://flip.id/tentang-flip> (accessed Jul. 20, 2023).
- [2] Darmawan, F. (2022). *TA: Analisis Kesuksesan Aplikasi Flip. Id Berdasarkan Persepsi Pengguna dengan Menggunakan Model DeLone and McLean* (Doctoral dissertation, Universitas Dinamika). [Online]. Available: <https://repository.dinamika.ac.id/id/eprint/6688/>
- [3] Cahyani, I. P. (2020). Membangun Engagement Melalui Platform Digital (Studi Kasus Flip sebagai Start-Up Fintech). *Ekspresi Dan Persepsi: Jurnal Ilmu Komunikasi*, 3(2), 76-87. DOI: 10.33822/jep.v3i2.1668.
- [4] Giovani, A. P., Ardiansyah, A., Haryanti, T., Kurniawati, L., & Gata, W. (2020). Analisis Sentimen Aplikasi Ruang Guru Di Twitter Menggunakan Algoritma Klasifikasi. *Jurnal Teknoinfo*, 14(2), 115-123. DOI: 10.33365/jti.v14i2.679.
- [5] Rahayu, S., Yumarlin, M. Z., Bororing, J. E., & Hadiyat, R. (2022). Implementasi Metode K-Nearest Neighbor (K-NN) untuk Analisis Sentimen Kepuasan Pengguna Aplikasi Teknologi Finansial FLIP. *Edumatic J. Pendidik. Inform*, 6(1), 98-106. DOI: 10.29408/edumatic.v6i1.5433.
- [6] Yustihan, S. R., Adikara, P. P., & Indriati, I. (2021). Analisis Sentimen berbasis Aspek terhadap Data Ulasan Rumah Makan menggunakan Metode Support Vector Machine (SVM). *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 5(3), 1017-1023.
- [7] Perdana, S. A. P., Aji, T. B., & Ferdiana, R. (2021). Aspect Category Classification dengan Pendekatan Machine Learning Menggunakan Dataset Bahasa Indonesia. *Jurnal Nasional Teknik Elektro dan Teknologi Informasi*, 10(3), 229-235. DOI: <https://doi.org/10.22146/jnteti.v10i3.1819>.
- [8] Sulistio, H. G., & Handojo, A. (2022). Aspect-Based Sentiment Analysis pada Ulasan ECommerce dengan Metode Support Vector Machine untuk Mendapatkan Informasi Sentimen dari Beberapa Aspek. *Jurnal Infra*, 10(2), 450-454.
- [9] Hamzah, F., Astuti, W., & Purbolaksono, M. D. (2022). Sentiment Analysis Pada Movie Review Menggunakan Feature Selection Chi Square Dan Support Vector Machine Classifier. *eProceedings of Engineering*, 9(3).

- [10] S. Prayogo, Y. Sibaroni, and S. Si, (2019). Aspect Based Sentiment Analysis Terhadap Ulasan Hotel Berbahasa Indonesia. Bandung.
- [11] ADMINLP2M, "Algoritma K-Nearest Neighbors (KNN) – Pengertian dan Penerapan.," *Lembaga Penelitian dan Pengabdian Masyarakat (LP2M) Universitas Medan Area*, Feb. 16, 2023. <https://lp2m.uma.ac.id/2023/02/16/algoritma-k-nearest-neighbors-knn-pengertian-dan-penerapan/> (accessed Aug. 07, 2023).
- [12] Amardita, R. S., Adiwijaya, A., & Purbolaksono, M. D. (2022). Analisis Sentimen terhadap Ulasan Paris Van Java Resort Lifestyle Place di Kota Bandung Menggunakan Algoritma KNN. *JURIKOM (Jurnal Riset Komputer)*, 9(1), 62-68. DOI: 10.30865/jurikom.v9i1.3793.
- [13] Fadhilah, J., Layyinna, C. A. A., Khatami, R., & Fitroh, F. (2021). Pemanfaatan Teknologi Digital Wallet Sebagai Solusi Alternatif Pembayaran Modern: Literature Review. *Journal of Computer Science and Engineering (JCSE)*, 2(2), 89-97. DOI: 10.36596/jcse.v2i2.219.
- [14] Arimuko, A., Wibawa, A. S. W., & Firmansyah, A. (2019). Analisis Perbandingan Penentuan Hiposentrum Menggunakan Metode Grid Search, Geiger, dan Random Search: Studi Kasus pada Letusan Gunung Sinabung 2017. *DIFFRACTION: Journal for Physics Education and Applied Physics*, 1(2), 22-28. DOI: <https://doi.org/10.37058/diffraction.v1i2.1290>
- [15] Mardiana, L., Kusnandar, D., & Satyahadewi, N. (2022). Analisis Diskriminan Dengan K Fold Cross Validation Untuk Klasifikasi Kualitas Air Di Kota Pontianak. *Bimaster: Buletin Ilmiah Matematika, Statistika dan Terapannya*, 11(1). DOI: <http://dx.doi.org/10.26418/bbimst.v11i1.51608>