

KLASIFIKASI ULASAN PENGGUNA APLIKASI ACCESS BY KAI BERBASIS ASPEK DENGAN ALGORITMA NAÏVE BAYES DAN SVM

Heru Indrawan, Bambang Irawan, Tati Suprpti

Program Studi Teknik Informatika, STMIK IKMI Cirebon

Jl. Perjuangan No.10 B, Kota Cirebon

heruindrawan04@gmail.com

ABSTRAK

Access by KAI adalah aplikasi resmi milik PT. Kereta Api Indonesia yang dirancang guna memenuhi kebutuhan para penumpangnya, baik yang melakukan perjalanan jarak jauh, menengah, maupun lokal. Meskipun Aplikasi *Access by KAI* masih memiliki beberapa kekurangan di beberapa aspek, oleh karena itu perbaikan harus terus dilakukan guna meningkatkan kualitasnya. Untuk menilai kualitas, kepuasan, dan kekurangan aplikasi ini, informasi dapat diperoleh dari ulasan di *Google Play Store* dan media sosial X. Penelitian ini bertujuan untuk mengidentifikasi sentimen terkait produk aplikasi *mobile* tersebut. dapat teridentifikasi sejumlah 6.000 ulasan yang sudah masuk. Hal ini bisa ditarik kesimpulan yang menegaskan kalau mayoritas pengguna sudah menyampaikan ulasan dengan lebih positif terkait aspek *satisfaction*, sedangkan untuk ulasan yang cenderung berlabel negative menjurus pada sejumlah aspek, seperti halnya *learnability*, *efficiency*, dan *errors*. Opini publik yang sudah dikelompokkan menjadi beberapa aspek terkait aplikasi *Access by KAI* dapat ditetapkan sebagai parameter kunci untuk menilai tingkat kepuasan masyarakat, sekaligus diharapkan menjadi dasar untuk melaksanakan evaluasi terhadap PT. Kereta Api Indonesia, khususnya pada aspek-aspek yang masih memiliki kekurangan dalam aplikasi *Access by KAI*. Perlu dipahami terkait penerapan dengan model CRISP-DM (*Cross Industry Standard Process for Data Mining*) dan komparasi hasil yang pijakannya dari metode klasifikasi *Naïve Bayes* dan *Support Vector Machine (SVM)*, sudah terbukti kalau model demikian menghadirkan kinerjanya paling maksimal. Jika diamati dari *mean* yang sudah melewati tahap uji di tiap aspeknya, yakni *Accuracy* besarnya 96.57%, *Precision* 94.52%, *Recall* 99.46%, dan *F-measure* 96.82%.

Kata Kunci: *Access by KAI, Analisis Sentimen, Aspect-Based, Naïve Bayes, Support Vector Machine.*

1. PENDAHULUAN

Perkembangan teknologi memberikan dampak besar pada layanan dalam suatu perusahaan. Perusahaan dapat memanfaatkan teknologi informasi untuk menyediakan layanan sekaligus informasi dengan lebih akurat dan cepat diterima. Dewasa ini, banyak perusahaan bersaing untuk menciptakan aplikasi *mobile* guna meningkatkan efisiensi dalam penyampaian informasi dan layanan [1]. Salah satu contoh perusahaan yang mengambil langkah tersebut adalah PT. Kereta Api Indonesia, seperti yang telah diketahui bersama berada di bawah naungan BUMN yang memiliki fokus atas tersedianya layanan angkutan untuk publik dengan basisnya yakni kereta api. Untuk meningkatkan pelayanan kepada pengguna kereta api di Indonesia, PT. KAI meluncurkan aplikasi *mobile* pada tahun 2014 yang sebelumnya dikenal sebagai *KAI Access*, dan sejak tanggal 10 Agustus 2023, aplikasi tersebut berganti menjadi *Access by KAI* [2]. *Mobile ticketing* adalah aplikasi seluler yang dapat digunakan untuk memesan dan membeli tiket [3]. Terdapat sejumlah aplikasi yang masuk pada kategorisasi ini, cakupannya cukup beragam, seperti Traveloka, kemudian ada juga yang namanya Tiket.com, lalu Pegipegi, hingga Tokopedia, yang mana aspek kajian di sini menjurus pada *Access by KAI*. Kendati demikian, masih disayangkan terkait kuantitas jumlah pengguna yang disediakan pihak PT. KAI ialah *Access by KAI* masih terkategori rendah dengan

dasarnya yakni ulasan yang ada di *Google Play Store* dan *platform* media sosial X.

Di lansir dari Kompas.com pembelian tiket melalui *platform Access by KAI* seringkali mengalami kendala teknis, terutama menjelang Hari Raya dan selama promo *Flash Sale*. Pengguna juga kerap menghadapi kesulitan dalam proses pengembalian tiket atau *refund* dana [4]. Realitas demikian menjadi satu dari banyaknya faktor yang menjadi efek pendorong pihak peneliti untuk melangsungkan pengkajian ilmiah secara lebih mendetail terkait *Access by KAI* sekaligus derajat rasa puas pelanggannya terkait aplikasi berbasis transportasi demikian. Perlu dipahami juga, kalau ada muatan dari segi yang krusial di mana butuh perhatian dalam urusan pengembangan pada aplikasi yang mengarah ke persepsi dari pihak pengguna atas sejumlah keluhan yang didasarkan pada mekanisme pemesanan tiket kereta api. Saat pengguna merasakan kepuasan dengan suatu pengadaan aplikasi, mereka akan mempunyai tendensi untuk senantiasa melibatkan penggunaannya.

Berlandaskan pada studi sebelumnya tentang metode klasifikasi dan penerapan Teknik SMOTE. Penelitian ini akan menyajikan perbandingan hasil antara metode klasifikasi *Naïve Bayes* dan *SVM*. Dengan menggunakan teknik SMOTE untuk mengatasi ketidakseimbangan kelas pada dataset ulasan *Access by KAI*, di mana ulasan negative lebih

banyak dibandingkan ulasan positive. Selain itu, dataset dalam penelitian ini dipecah menjadi beberapa aspek, yaitu *learnability*, *efficiency*, *errors*, dan *satisfaction*. Metode Naïve Bayes dan SVM, adalah teknik populer dalam klasifikasi teks, di mana metode-metode tersebut telah belajar dari contoh-contoh sebelumnya guna menghadirkan hasil secara lebih maksimal. Temuan pada riset ini tentu diharapkan dapat menjadi jalan bagi kelangsungan proses identifikasi terkait algoritma yang menjurus pada derajat *accuracy*, lalu menjurus ke *recall*, tidak ketinggalan perwujudan dari *precision*, dan yang terakhir kajian perihal *F-measure* terbaik di dalam kategorisasi ulasan pengguna *Access by KAI*.

Temuan riset ini menjadi pengharapan bersama untuk senantiasa bisa menghadirkan kontribusi dengan sifat yang signifikan dengan sasarannya yakni pada ranah pengembangan Aplikasi *Access by KAI* dan manajemen kualitas layanan PT. Kereta Api Indonesia. Hasil klasifikasi ini dapat digunakan sebagai panduan dalam meningkatkan aspek-aspek tertentu yang dinilai negatif oleh pengguna, seperti pemahaman, efisiensi, dan penanganan kesalahan. Dengan memahami sentimen pengguna melalui algoritma Naïve Bayes dan SVM, perusahaan dapat merancang strategi perbaikan yang lebih tepat sasaran, meningkatkan kepuasan pengguna, dan menjaga reputasi Aplikasi *Access by KAI* di mata publik.

2. TINJAUAN PUSTAKA

Penelitian sebelumnya yang dilakukan oleh Rahayu, dkk. (2022) membahas perbandingan antara Algoritma Naïve Bayes dan *Support Vector Machine* yang bertalian dengan mekanisme penganalisisan atas sentiment pihak pengguna dari *Spotify*. Berpijak pada hasilnya, bisa nampak kalau sebagian besar ulasan cenderung positif. Temuan pada riset menegaskan kalau akurasi dari algoritma SVM mencapai persentase 84%, sedangkan untuk akurasi algoritma Naïve Bayes lebih tinggi ketimbang *Support Vector Machine*, yakni mencapai 86,4% [5].

Nikmatul, dkk. (2019) menggunakan teknik SMOTE guna mengatasi ketidakseimbangan kelas dalam proses klasifikasi objek berita online dengan memanfaatkan algoritma k-NN. Sementara data yang dilibatkan asalnya dari hasil penghimpunan informasi pada salah satu portal berita terkemuka, yakni *kompas.com*, dan mengungkapkan adanya ketidakseimbangan antara jumlah berita subjektif dan berita objektif. Ketidakseimbangan tersebut berpotensi memengaruhi performa klasifikasi secara keseluruhan. Dalam konteks penelitian ini, nilai k variasi yang diujikan adalah 1, 3, 5, 7, dan 9. Penerapan SMOTE berhasil memberikan peningkatan yang signifikan pada akurasi algoritma. Hasil terbaik dari klasifikasi SMOTE dengan k-NN tercapai pada nilai $k = 1$, dengan tingkat akurasi mencapai 87,50%, *presisi* 0,80, *recall* 0,98, dan *F-measure* 0,88 [6].

Berlanjut ke riset yang berhasil dilangsungkan Iskandar, dkk. (2021) dengan kajian yang

mengkomparasikan penganalisisan sentiment gadget dengan basisnya mengacu pada metode Naïve Bayes, lalu SVM, dan terakhir yakni k-NN. Lebih lanjut soal sumber data yang sengaja dilibatkan pada riset asal-muasalnya yakni dari komentar di platform *YouTube* dengan keterkaitannya yang masif pada penggunaan gadget Samsung *Galaxy Z Flip 3*. Temuan pada riset tersebut menegaskan kalau untuk model klasifikasi SVM menghadirkan hasil paling baik. Secara menyeluruh sendiri, akurasi untuk SVM meraih persentase angka yakni 96,43% teruntuk ke-4 aspek pokok yang dilibatkan bagi kepentingan pengujian, dengan cakupannya yakni desain 94,40%, selanjutnya berlanjut ke harga 97,44%, lalu kaitannya dengan spesifikasi 96,22%, dan terakhir kajian atas citra mereknya yakni 97,63% [7].

Riset yang digagas oleh Salsabila, dkk. (2021) melangsungkan pengkajian terkait ulasan produk berbasis kecantikan di komunitas *Female Daily* yang mana menitikberatkan pada mekanisme klasifikasi SVM. Lebih lanjut soal data set yang sengaja dilibatkan ada sangkutpautnya dengan kategori produk, misalnya *serum & essence*, *toner*, *scrub & exfoliating*, dan *sunscreen*. Adapun temuan terkait kategorisasinya yang menjurus pada eksistensi produk mencakup 3 aspek yang sengaja dilibatkan, yakni terkait dengan harga, kemudian mengarah ke aspek kemasan, lalu yang terakhir perihal aroma. Derajat akurasi dari basis optimal yang didapatkannya yakni meraih persentase 93% atas akurasi dari segi harga. Dan soal kemasan mendapatkan persentase 92%, sedangkan bagian aroma perolehannya yakni 86% [8].

2.1. Data Mining

Perihal data mining dapat diberikan label definisi sebagai serangkaian proses yang ada sangkut pautnya dengan praktik satu atau bahkan lebih teknik *machine learning* yang diperuntukkan secara otomatis dalam proses menganalisis sekaligus pengeksktraksian ragam pengetahuan. Adapun tahapan dari data mining memiliki tujuan spesifik untuk melangsungkan skema identifikasi terkait informasi pada representasi himpunan data yang didalamnya mengikutsertakan kebermanfaatan atas teknik ataupun algoritma secara spesifik [9].

2.2. Sentiment Analysis (SA)

Topik yang mengarah pada analisis sentimen dapat diberikan label definisi sebagai serangkaian proses otomatis Dengan pemahaman dan tata kelola data teks yang memiliki tujuan mendapatkan informasi terkait sentimen yang ada pada suatu kalimat atau perwujudan teks dengan representasi berupa pendapat. Hal ini dilangsungkan untuk kepentingan evaluasi atas perspektif atau bisa juga dikaitkan dengan opini yang memuat pengadaan suatu teks dengan keterkaitan objek atau persoalan yang diusung apakah sifatnya positif atau justru sebaliknya ialah negatif. Adapun komponen analisis sentimen mencakup sejumlah hal yang bersifat cukup kompleks mulai dari pemrosesan

bahasa dengan alami, lalu berlanjut pada penganalisisan teks dan yang ketiga terkait komputasi dari perspektif linguistik yang diperuntukkan bagi identifikasi sentimen atas mekanisme dokumen yang ada [10].

2.3. Aspect-based sentiment analysis (ABSA)

Poin yang mengacu pada analisis sentimen berpijak pada aspek yang bertalian dengan sentimen yang terperinci tugas klasifikasi yang bertujuan untuk mengidentifikasi polaritas sentimen (positif, netral, negatif) dari aspek dalam kalimat tertentu [11].

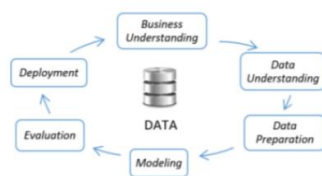
2.4. Klasifikasi

Klasifikasi merupakan tahapan dalam melangsungkan pencarian atas keberfungsian atau eksistensi model yang bisa menerangkan pun membedakan konsep dan kelas pada data. Sedangkan tujuannya sendiri yakni melangsungkan prediksi dari kelas objek yang kelasnya belum teridentifikasi dengan jelas [12].

2.5. Confus Matrix

Bahasan yang mengarah pada *Confusion Matrix* merujuk ke suatu teknik yang umumnya melibatkan demi kepentingan kalkulasi atas derajat akurasi. Saat menguji akurasi hasil pencarian, nilai *recall*, *precision*, *akurasi*, dan tingkat kesalahan akan dievaluasi atau *f-measure* [13].

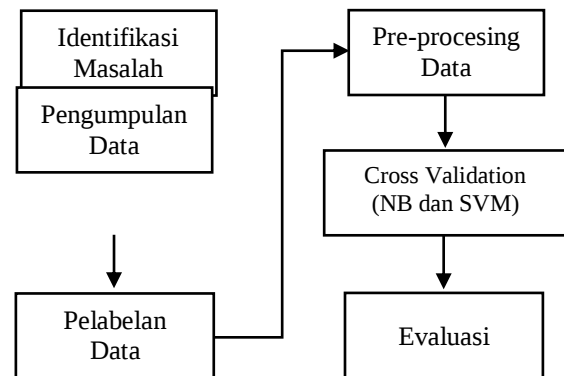
3. METODE PENELITIAN



Gambar 1. Model CRISP-DM

Riset ini menitikberatkan pada metode khusus bertajuk CRISP-DM sebagaimana yang terpapar melalui gambar 1. CRISP-DM (*Cross-Industry Standard Process for Data Mining*) di mana hal tersebut merujuk pada serangkaian metode yang bisa dilibatkan demi kepentingan strategi guna merampungkan sejumlah problematik *general* sekaligus terkait dengan *metodologi* guna melangsungkan penetapan atas standarisasi *general* yang mengarah ke data mining yang dikaitkan [9]. CRISP-DM melangsungkan pola pembentukan atas standarisasi proses ekstraksi menjurus pada informasi yang mengarah ke perwujudan data yang hendak diaplikasikan menjadi pondasi untuk strategi merampungkan problematik *general* dalam konteks bisnis atau unit kajian. Metode CRISP-DM mencakup sejumlah fase, termasuk *Business Understanding*, selanjutnya terkait dengann *Data Understanding*, lalu mengarah ke unsur *Data Preparation*, berlanjut ke bagian *Modelling*, dan terakhir yakni *Deployment* [14].

Perlu dipahami kalau riset ini hanya menekankan pada keterlibatan tahapan evaluasi dan tidak dilanjutkan ke tahapan yang lebih kompleks, yakni *Deployment*, yang lekat dengan langkah implementasi pada *tools*.



Gambar 2. Tahapan metode penelitian

Pada riset ini, pihak peneliti menerapkan Model CRISP-DM, sehingga langkah-langkah penelitiannya disesuaikan sesuai dengan yang ditunjukkan dalam Gambar 2. Tahapan penelitian melibatkan identifikasi masalah, pengumpulan data ulasan *Access by KAI* dari dataset yang berasal dari *Google Play Store* dan Media sosial X, pelabelan data, kemudian terkait dengan *pre-processing*, lalu menjurus ke aspek *cross-validation* menggunakan algoritma Naïve Bayes (NB) dan SVM, dan terakhir yakni tahap evaluasi.

3.1. Business Understanding (Pemahaman Bisnis)

Tahapan bisnis understanding diupayakan dengan keterkaitan pemahaman atas objek atau problematik yang diusung selama riset berlangsung. Persoalan yang diupayakan penganalisisan mencakup sejauh mana kajian terkait *accuracy*, selanjutnya *presisi*, lalu *recall*, dan terakhir yakni *F-measure* yang yang dihasilkan dari implementasi algoritma Naïve Bayes dan SVM. Selanjutnya terkait dengan dataset yang dilibatkan didapatkan melalui ulasan pihak pengguna yang terjali pada aplikasi bertajuk *Access by KAI*.

3.2. Data Understanding (Pemahaman Data)

Mengarah pada tahap pemahaman data diupayakan sejumlah proses yang mengacu pada pengambilan data dengan sifat yang masih mentah untuk kemudian disejajarkan dengan atribut yang memang menjadi kebutuhan. Perolehan data berpijak pada ulasan *Google Play Store* dan media sosial X. Data dikumpulkan dari Januari 2023 hingga November 2023, mencakup sebanyak 6.000 data yang terkait dengan ulasan di dalamnya. Selanjutnya terkait ulasan yang dipaparkan pihak pengguna tidak relevan dengan riset, misalnya terkait pertanyaan promo, atau nama dari versi aplikasi tersebut, dan hal-hal yang tidak perlu untuk dilibatkan sebagai perwujudan dari dataset. Selepas melewati proses seleksi ditemukan kisaran 1431 ulasan yang berhasil dipilih. Kemudian data set tersebut dilibatkan untuk kepentingan

pemberian label dengan manual yang selanjutnya diikuti proses pengidentifikasian atas aspeknya. Sementara ada penetapan terkait 4 kategori yang menjurus pada konteks *Learnability* (Mudah Dipelajari), selanjutnya terkait dengan *Efficiency* (Efisien), lalu persoalan *Errors* (Pencegahan Kesalahan), dan terakhir yakni *Satisfaction* (Kepuasan Pengguna). Perlu dipahami kalau aspek tersebut didasarkan atas penganalisisan data terkait ulasan yang ada di aplikasi *Access by KAI*. Aspek *Learnability* (Le), Memaparkan terkait seberapa mudahnya para pengguna bisa merampungkan sejumlah tugas dasar saat pertama kali mengakses dan menggunakan aplikasi tersebut. Aspek *Efficiency* (Ef) menerangkan perihal efektivitas aplikasi dan layanan *Access by KAI*. Aspek *Errors* (Er) menerangkan terkait kuantitas error yang acap kali dialami pihak pengguna pada saat menggunakan aplikasi tersebut. Selanjutnya soal aspek *Satisfaction* (Sa) mengacu pada derajat rasa puas pihak pengguna ketika mengakses dan menggunakan aplikasi berbasis *Access by KAI*. Berdasarkan Tabel 1. Aspek yang dihasilkan kemudian diklasifikasikan pada dua kategori pokok yakni positif dan juga negatif. Pelabelan positif ditandai dengan angka 1 sementara negatif dengan - 1. Sementara untuk ulasan yang tidak mencakup ruang lingkup atas riset yang dikaji diberikan label 0.

Tabel 1. Labeling dataset

No	Ulasan	Le	Ef	Er	Sa
1.	Mudah Praktis saat digunakan	1	0	0	0
2.	Buat batalin tiket aja lama bngt di konfirmasi nya sudah dua minggu	0	-1	0	0
3.	Parah banyak errornya mending versi yang lama	0	0	-1	0
4.	Sangat menyenangkan membantu dalam perjalanan saya untuk rutinitas luar kota	0	0	0	1

3.3. Data Preparation (Persiapan Data)

Fase penyiapan data, sering disebut sebagai tahap preprocessing, merujuk pada langkah-langkah untuk memastikan data siap pakai dan bersih dalam konteks penelitian akademis. Beberapa proses *preprocessing* melibatkan *transformasi* Kasus, di mana huruf besar diubah menjadi huruf kecil. *Tokenisasi*, yaitu menghapus tanda baca, simbol, karakter, dan tanda baca yang dianggap tidak relevan. Filter Token (Berdasarkan Panjang) digunakan untuk menetapkan rentang panjang minimum dan maksimum dari ulasan. Penghapusan Stopwords melibatkan eliminasi kata-kata yang dianggap tidak esensial. Di sisi lain, stemming adalah proses penyederhanaan kata-kata dengan menghapus imbuhan sehingga hanya menyisakan kata dasar.

3.4. Modelling (Pemodelan)

Dalam tahap pemodelan ini, suatu model dikembangkan dengan menerapkan teknik klasifikasi pada dataset ulasan yang telah melalui proses preprocessing sebelumnya. Algoritma klasifikasi yang diterapkan mencakup Naïve Bayes dan SVM, dengan

penggunaan metode SMOTE untuk mengatasi ketidakseimbangan kelas. Penelitian ini menggunakan RapidMiner versi 10.2 sebagai alat bantu untuk menjalankan proses analisis dan pengembangan model.

3.5. Evaluation (Evaluasi)

Pada tahap evaluasi ini, metode klasifikasi dievaluasi dengan mengukur kinerjanya menggunakan confusion matrix terhadap algoritma Naïve Bayes dan SVM. *Confusion matrix* digunakan sebagai alat untuk menghitung akurasi, dan dalam pengujian keakuratan, nilai *accuracy* (formula 1), *precision* (formula 2), *recall* (formula 3), dan *f-measure* (formula 4).

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

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

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

$$F - measure = 2x \frac{precision \times recall}{precision+recall} \quad (4)$$

Ukuran evaluasi model klasifikasi melibatkan konsep True Positive (TP), yang mencakup jumlah data dengan nilai positif yang diprediksi dengan benar sebagai positif. Selain itu, True Negative (TN) menggambarkan jumlah data yang memiliki nilai negatif dan diprediksi dengan benar sebagai negatif. Keberadaan False Positive (FP) menunjukkan jumlah data dengan nilai negatif yang, namun, diprediksi sebagai positif, mencerminkan kesalahan di mana model salah mengidentifikasi kejadian negatif. Di sisi lain, False Negative (FN) mewakili kumpulan data yang memiliki nilai positif, tetapi diprediksi sebagai negatif, mencerminkan kesalahan di mana model salah mengidentifikasi kejadian positif sebagai negatif.

4. HASIL DAN PEMBAHASAN

Dalam bab ini, akan diuraikan hasil dan pembahasan mengenai kinerja algoritma Naïve Bayes dan SVM dengan penerapan metode *Cross-Industry Standard Process for Data Mining* (CRISP-DM) yang digunakan dalam kerangka penelitian ini. Data yang digunakan diperoleh dari ulasan *Access by KAI* di *Google Play Store* dan media sosial X.

Tabel 2. Dataset berdasarkan aspek

Sentimen	Le	Ef	Er	Sa
Positif	21	20	19	137
Negatif	242	221	311	460
Total	263	241	330	597

Pada Tabel 2. merupakan hasil pelabelan dilakukan secara manual oleh Peneliti dengan bantuan Guru Bahasa Indonesia dari SMK NU Cantigi, yaitu Ibu Tarniah S.Pd. Pelabelan menggunakan label positif dan negatif. Setelah proses seleksi, sebanyak 1.431 ulasan diperoleh dan dikelompokkan berdasarkan aspek *Learnability* (Le), *Efficiency* (Ef), *Errors* (Er), dan *Satisfaction* (Sa).

4.1. Pre-processing

Sebelum digunakan untuk klasifikasi data, langkah awal yang dilakukan adalah proses pre-processing. Penelitian ini mengimplementasikan beberapa tahap pre-processing sebagai langkah persiapan data, yakni *Transform Case*, *Tokenizing*, *Filter Token (By Length)*, *Stopwords*, dan *Stemming*.

Tabel 3. Contoh tahap *transform case*

Sebelum	Sesudah
Bagusan Aplikasi lama mao pesan tiket mlh ga bisa jadi nyesel diperbarui	bagusan aplikasi lama mao pesan tiket mlh ga bisa jadi nyesel diperbarui

Pada Tabel 3. merupakan contoh hasil transformasi dari tahap *transform case*, teks ulasan diubah menjadi huruf kecil untuk memudahkan proses pengolahan.

Tabel 4. Contoh tahap *tokenizing*

Sebelum	Sesudah
bagusan aplikasi lama mao pesan tiket mlh ga bisa jadi nyesel diperbarui	'bagusan' 'aplikasi' 'lama' 'mao' 'pesan' 'tiket' 'mlh' 'ga' 'bisa' 'jadi' 'nyesel' 'diperbarui'

Pada Tabel 4. merupakan contoh hasil transformasi dari tahap *Tokenizing*, teks ulasan diubah dari bentuk kalimat menjadi Kumpulan kata.

Tabel 5. Contoh tahap *token (by length)*

Sebelum	Sesudah
'bagusan' 'aplikasi' 'lama' 'mao' 'pesan' 'tiket' 'mlh' 'ga' 'bisa' 'jadi' 'nyesel' 'diperbarui'	'bagusan' 'aplikasi' 'lama' 'pesan' 'tiket' 'mlh' 'bisa' 'jadi' 'nyesel' 'diperbarui'

Pada Tabel 5. diberikan ilustrasi tentang implementasi tahap *Filter Token (By Length)*. Langkah ini bertujuan untuk mengeliminasi kata-kata yang dianggap tidak relevan setelah proses *Tokenizing*. Dalam pelaksanaannya, kata-kata dengan panjang tertentu akan dihapus. Dalam konteks ulasan, kata-kata yang memiliki panjang kurang dari 4 kata atau lebih dari 30 kata dihapus. Sebagai contoh, istilah seperti "yg," "tdk," "jd," "ga," "ane," "mau," dan "gan" dianggap tidak memiliki makna khusus ketika dipisahkan dari kata-kata lainnya dan tidak berkaitan dengan kata sifat yang terkait dengan sentimen.

Tabel 6. Contoh tahap *stopwords*

Sebelum	Sesudah
'bagusan' 'aplikasi' 'lama' 'pesan' 'tiket' 'mlh' 'bisa' 'jadi' 'nyesel' 'diperbarui'	'bagusan' 'aplikasi' 'pesan' 'tiket' 'nyesel' 'diperbarui'

Pada Tabel 6. merupakan contoh hasil transformasi dari tahap *Stopwords*. Pada tahap stopwords, kata-kata yang tidak penting akan

dihilangkan untuk menghasilkan teks yang lebih bermakna dan relevan dengan klasifikasi sentiment.

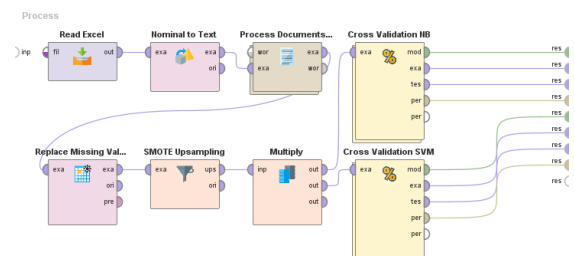
Tabel 7. Contoh tahap *stemming*

Sebelum	Sesudah
'bagusan' 'aplikasi' 'pesan' 'tiket' 'nyesel' 'diperbarui'	'bagus' 'aplikasi' 'pesan' 'tiket' 'menyesel' 'baru'

Pada Tabel 7. merupakan contoh hasil transformasi dari tahap *stemming*. Pada Tahap *Stemming*, kata-kata diubah menjadi bentuk dasarnya agar lebih mudah diproses.

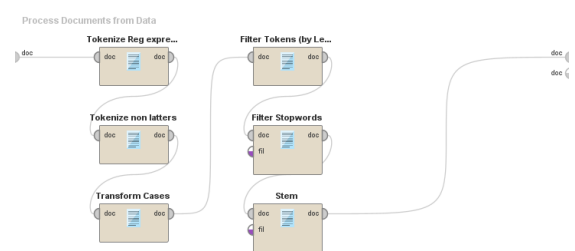
4.2. Pembuatan Model Klasifikasi

Langkah pembuatan model klasifikasi merupakan fase kunci dalam upaya membangun model yang dapat digunakan untuk mengklasifikasikan data ulasan setelah melalui tahap pra-pemrosesan. Dalam penelitian ini, diterapkan tiga algoritma klasifikasi, yaitu Naive Bayes dan SVM. Proses tersebut didukung oleh penggunaan alat bantu RapidMiner versi 10.2.



Gambar 3. Proses pembuatan model klasifikasi

Gambar 3. menunjukkan alur proses pemodelan. Setelah komentar diberi label secara manual dalam format *Microsoft Excel*. Setelah itu, data akan diproses oleh operator *Nominal to Text* untuk mengubah atribut yang bernilai nominal menjadi *string*.



Gambar 4. Proses dalam operator process document

Proses dilanjutkan dengan operator *Process Documents* yang ditunjukkan pada gambar 4. atau yang disebut dengan tahap data *preparation (pre-processing)*. Dimulai dari tahap *Tokenizing*, *transform case*, *Filter Token (By Length)*, *Stopwords*, dan *Stemming*.

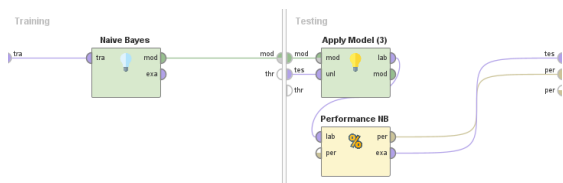
Setelah data menjalani proses oleh operator *Process Document*, langkah selanjutnya melibatkan operator *Replace Missing Values* untuk mengisi nilai yang hilang dalam atribut. Proses ini dirancang untuk

menangani nilai yang kosong dalam dataset. Selanjutnya, untuk mengatasi ketidakseimbangan kelas pada dataset ulasan Access by KAI, di mana ulasan negative lebih banyak dibandingkan ulasan positive. Penggunaan operator SMOTE *Up Sampling* bertujuan untuk menyelaraskan dataset sehingga kelasnya menjadi seimbang, dan hal ini akan diperkuat dengan penggunaan operator *Multiply*. Operator *Multiply* memungkinkan pengguna untuk menerapkan dua metode *Cross Validation* secara bersamaan, mendukung pendekatan yang lebih holistik dalam evaluasi model.

Tabel 8. Proporsi data menggunakan SMOTE

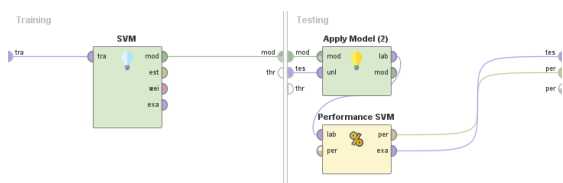
Sentimen	Le	Ef	Er	Sa
Positif	242	221	311	460
Negatif	242	221	311	460

Dari Tabel 8, dapat dilihat bahwa Operator SMOTE *Up Sampling* ditambahkan untuk mengatasi ketidak seimbangan data sentiment positif maupun negatif.



Gambar 5. Desain evaluasi Naive Bayes

Pada Gambar 5. Merupakan Desain Evaluasi Naive Bayes menggunakan metode *Cross Validation* dengan 10 lipatan (*folds*) melibatkan pembagian data menjadi 10 bagian, di mana setiap *iterasi* 9 bagian digunakan sebagai data pelatihan (*training*) dan 1 bagian digunakan sebagai data pengujian (*testing*). Proses ini diulangi sebanyak 10 kali, sehingga setiap bagian menjadi set pengujian satu kali. Hasil akhir evaluasi performa Naive Bayes diperoleh dengan mengambil nilai rata-rata dari semua iterasi pengujian tersebut. Pendekatan ini memungkinkan evaluasi yang lebih *robust* terhadap performa model Naive Bayes di berbagai bagian dataset.



Gambar 6. Desain evaluasi SVM

Gambar 6. menampilkan Desain Evaluasi Support Vector Machine (SVM) menggunakan metode *Cross Validation* dengan 10 lipatan (*folds*). Dalam proses ini, data dibagi menjadi 10 bagian, di mana setiap iterasi melibatkan 9 bagian sebagai data pelatihan dan 1 bagian sebagai data pengujian. Iterasi ini diulangi sebanyak 10 kali, sehingga setiap bagian dataset berperan sebagai set pengujian satu kali.

Evaluasi akhir performa SVM dihasilkan dengan mengambil nilai rata-rata dari semua iterasi pengujian. Pendekatan ini memberikan evaluasi yang lebih andal terhadap performa model SVM di berbagai bagian dataset.

4.3. Evaluasi Model Klasifikasi

Performa algoritma Naive Bayes dan SVM dalam mengklasifikasikan sentimen ulasan pengguna Access by KAI akan dianalisis. Analisis ini mencakup faktor-faktor yang mempengaruhi akurasi algoritma. Performa algoritma dapat dilihat dari nilai *accuracy* (ac), *precision* (pr), *recall* (re), dan *F-measure* (f-m) yang dihasilkan confusion matrix dari setiap aspek.

Tabel 9. Hasil pengujian aspek *learnability*

Algoritma	Ac	Pr	Re	F-m
NB	99,17%	98,37%	100,00%	99,18%
SVM	89,83%	83,10%	100,00%	90,77%

Berdasarkan data pada Tabel 9, terlihat bahwa hasil pengujian pada aspek *Learnability*. Algoritma Naive Bayes lebih unggul dibandingkan dengan model klasifikasi SVM. Empat hasil tertinggi dari algoritma Naive Bayes mencakup *Accuracy* sebesar 99.17%, *Precision* sebesar 98.37%, *Recall* sebesar 100.00%, dan aspek *F-measure* sebesar 99.18%. Rata-rata dari masing-masing model klasifikasi adalah 99.18% untuk Naive Bayes dan 90.93% untuk SVM.

Tabel 10. Hasil pengujian aspek *efficiency*

Algoritma	Ac	Pr	Re	F-m
NB	99,77%	99,55%	100,00%	99,77%
SVM	84,77%	76,66%	100,00%	86,79%

Hasil pengujian pada aspek *Efficiency* ditunjukkan pada Tabel 10. urutan hasil tertinggi dalam pada Aspek *Efficiency* yaitu algoritma Naive Bayes pada *Accuracy* sebesar 99.77%, *Precision* sebesar 99.55%, *Recall* sebesar 100.00%, dan *F-measure* sebesar 99.77%. Hasil rata-rata dari masing-masing klasifikasi yaitu Naive Bayes sebesar 99.77% dan SVM sebesar 87.05%. Pada hasil precision, model algoritma Naive Bayes memiliki nilai rata-rata yang unggul.

Tabel 11. Hasil pengujian aspek *errors*

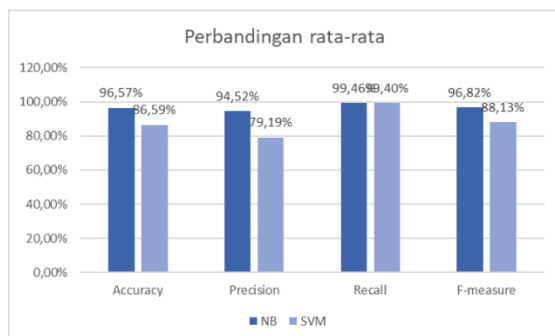
Algoritma	Ac	Pr	Re	F-m
NB	98,87%	97,79%	100,00%	98,88%
SVM	86,45%	78,68%	100,00%	88,07%

Berdasarkan data pada Tabel 11, terlihat bahwa hasil pengujian pada aspek *errors*. Algoritma Naive Bayes lebih unggul dibandingkan dengan model klasifikasi SVM. Empat hasil tertinggi dari algoritma Naive Bayes mencakup *Accuracy* sebesar 98.87%, *Precision* sebesar 97.79%, *Recall* sebesar 100.00%, dan aspek *F-measure* sebesar 98.88%. Rata-rata dari masing-masing model klasifikasi adalah 98.89% untuk Naive Bayes dan 88.30% untuk SVM.

Tabel 12. Hasil pengujian aspek *satisfaction*

Algoritma	Ac	Pr	Re	F-m
NB	88,45%	82,39%	97,82%	89,44%
SVM	85,30%	78,32%	97,60%	86,90%

Data pada Tabel 12. terlihat bahwa hasil pengujian pada aspek *satisfaction*. Algoritma Naïve Bayes lebih unggul dibandingkan dengan model klasifikasi SVM. Empat hasil tertinggi dari algoritma Naïve Bayes mencakup *Accuracy* sebesar 88.45%, *Precision* sebesar 82.39%, *Recall* sebesar 97.82%, dan aspek *F-measure* sebesar 89.44%. Rata-rata dari masing-masing model klasifikasi adalah 89.53% untuk Naïve Bayes dan 87.03% untuk SVM.



Gambar 7. Perbandingan rata-rata model klasifikasi

Hasil yang mengarah ke komparasi terkait rerata selepas proyeksi dilangsungkan dalam perwujudan grafik bisa diamati dengan gamblang melalui Gambar 7. Berlandaskan pada komparasi rerata atas hasil dari model klasifikasi yang terlibat, yakni Naïve Bayes menegaskan rerata tertinggi dikomparasikan dengan SVM. Hal ini terlihat dari nilai rata-rata *accuracy* sebesar 96,57%, *precision* sebesar 94,52%, *recall* sebesar 99,46%, dan *f-measure* sebesar 96,82%.

5. KESIMPULAN DAN SARAN

Berlandaskan pada hasil penganalisan sekaligus pengujian ulasan pengguna *Access by KAI* dari *Google Play Store* dan media sosial X, umumnya, sentimen positif terhadap aplikasi *Access by KAI* berkaitan dengan kemampuan pembelian tiket tanpa perlu pergi ke stasiun. Di sisi lain, sentimen negatif terhadap aplikasi ini terkait dengan munculnya *error* setelah pembaruan, kesulitan dalam registrasi awal, serta kesulitan mendapatkan tiket kereta pada periode Lebaran maupun selama promo. Dalam membandingkan metode Naive Bayes dan SVM, terungkap bahwa model klasifikasi Naive Bayes memberikan hasil terbaik. Dilihat dari hasil rata-rata pengujian pada setiap aspek, yakni *Accuracy* sebesar 96.57%, *Precision* sebesar 94.52%, *Recall* sebesar 99.46%, dan *F-measure* sebesar 96.82%. temuan pada riset juga menghadirkan saran yang bisa diadopsi demi kepentingan riset lanjutan, misalnya melakukan pengujian dengan melibatkan algoritma yang terhidang secara berbeda dan tidak terbatas pada penganalisan sentimen ulasan *Access by KAI*,

melainkan juga dengan topik yang berbeda. Selain itu, perlu dilakukan analisis mendalam terkait dataset yang digunakan, termasuk penerapan pelabelan pada setiap data. Dalam tahap *pre-processing*, dianjurkan menambahkan proses-proses tambahan guna meningkatkan hasil klasifikasi dan prediksi.

DAFTAR PUSTAKA

- [1] n. P. Ayu wangi diantini, s. Sukidin, and w. Hartanto, "efektivitas penerapan mobile application 'kai access' oleh konsumen di pt. Kereta api indonesia persero daerah operasi 9 stasiun jember," *jurnal pendidikan ekonomi: jurnal ilmiah ilmu pendidikan, ilmu ekonomi dan ilmu sosial*, vol. 13, no. 2, p. 132, sep. 2019, doi: 10.19184/jpe.v13i2.11477.
- [2] i. E. P. Y. C. Asia sanjaya, "kai access berubah nama jadi access by kai mulai 10 agustus 2023, apa saja fiturnya?komentar: baca berita tanpa iklan. Gabung Kompas.com+ Kompas.com tren kai access berubah nama jadi access by kai mulai 10 agustus 2023, apa saja fiturnya? Artikel ini telah tayang di Kompas.com dengan judul 'kai access berubah nama jadi access by kai mulai 10 agustus 2023, apa saja fiturnya?'," Kompas.com.
- [3] g. Radiena and a. Nugroho, "analisis sentimen berbasis aspek pada ulasan aplikasi kai access menggunakan metode support vector machine," 2023.
- [4] y. C. A. S. Farid firdaus, "aplikasi access by kai eror saat 'flash sale' tiket kereta rp 78.000, ini kata kai artikel ini telah tayang di Kompas.com dengan judul 'aplikasi access by kai eror saat 'flash sale' tiket kereta rp 78.000, ini kata kai,'" Kompas.com.
- [5] a. S. Rahayu, a. Fauzi, and r. Rahmat, "komparasi algoritma naïve bayes dan support vector machine (svm) pada analisis sentimen spotify," *jurnal sistem komputer dan informatika (json)*, vol. 4, no. 2, p. 349, dec. 2022, doi: 10.30865/json.v4i2.5398.
- [6] s. Keputusan dirjen penguatan riset dan pengembangan ristek dikti, a. Nikmatul kasanah, u. Pujiyanto, t. Elektro, f. Teknik, and u. Negeri malang, "terakreditasi sinta peringkat 2 penerapan teknik smote untuk mengatasi imbalance class dalam klasifikasi objektivitas berita online menggunakan algoritma knn," *masa berlaku mulai*, vol. 1, no. 3, pp. 196–201, 2017.
- [7] j. W. Iskandar and y. Nataliani, "perbandingan naïve bayes, svm, dan k-nn untuk analisis sentimen gadget berbasis aspek," *jurnal resti (rekayasa sistem dan teknologi informasi)*, vol. 5, no. 6, pp. 1120–1126, dec. 2021, doi: 10.29207/resti.v5i6.3588.
- [8] irbah salsabila and yuliant sibaroni, "multi aspect sentiment of beauty product reviews using svm and semantic similarity," *jurnal resti (rekayasa sistem dan teknologi informasi)*, vol. 5, no. 3, pp.

- 520–526, jun. 2021, doi: 10.29207/resti.v5i3.3078.
- [9] a. Dwiki, a. Putra, and s. Juanita, “analisis sentimen pada ulasan pengguna aplikasi bibit dan bareksa dengan algoritma knn,” vol. 8, no. 2, 2021, [online]. Available: <http://jurnal.mdp.ac.id>
- [10] a. Imron, “analisis sentimen terhadap tempat wisata di kabupaten rembang menggunakan metode naive bayes classifier.”
- [11] l. Xu and w. Wang, “improving aspect-based sentiment analysis with contrastive learning,” *natural language processing journal*, vol. 3, p. 100009, jun. 2023, doi: 10.1016/j.nlp.2023.100009.
- [12] i. Budiman and r. Ramadina, “penerapan fungsi data mining klasifikasi untuk prediksi masa studi mahasiswa tepat waktu pada sistem informasi akademik perguruan tinggi,” *ijccs*, vol. X, no.x, no. 1, pp. 1–5.
- [13] v. Amrizal, “penerapan metode term frequency inverse document frequency (tf-idf) dan cosine similarity pada sistem temu kembali informasi untuk mengetahui syarah hadits berbasis web (studi kasus: hadits shahih bukhari-muslim),” *jurnal teknik informatika*, vol. 11, no. 2, pp. 149–164, nov. 2018, doi: 10.15408/jti.v11i2.8623.
- [14] muhammad andi yudha, “crisp-dm, pendekatan proses dalam data mining,” medium.