

# Perbandingan Kinerja SVM, Random Forest dan XGBoost pada Aplikasi Access by KAI Menggunakan ADASYN

Nadia Epriyanti, Allsela Meiriza\*, Dinna Yunika Hardiyanti

Fakultas Ilmu Komputer, Sistem Informasi, Universitas Sriwijaya, Palembang, Indonesia  
Email: '09031182227017@student.unsri.ac.id, <sup>2\*</sup>allsela@unsri.ac.id, <sup>3</sup>dinna.yunika@gmail.com  
Email Penulis Korespondensi: allsela@unsri.ac.id  
Submitted 15-10-2025; Accepted 30-10-2025; Published 31-10-2025

## Abstrak

Pesatnya perkembangan aplikasi digital mendorong kebutuhan untuk memahami persepsi pengguna secara mendalam, salah satunya melalui analisis sentimen terhadap ulasan pengguna. Dalam praktiknya, proses klasifikasi sentimen sering kali menghadapi kendala distribusi data yang tidak seimbang, khususnya ketika jumlah ulasan dengan kategori netral jauh tertinggal dibandingkan dua kategori lainnya, yaitu positif dan negatif. Ketimpangan ini dapat memengaruhi kinerja model dalam mengidentifikasi semua kategori sentimen secara optimal. Studi ini dilakukan untuk menguji perbandingan hasil kinerja tiga algoritma *machine learning*, yaitu *Support Vector Machine* (SVM), *Random Forest* (RF), dan *Extreme Gradient Boosting* (XGBoost), dengan penerapan teknik *Adaptive Synthetic Sampling* (ADASYN) untuk mengatasi ketidakseimbangan distribusi kelas. Penelitian ini berbeda dengan penelitian terdahulu yang sejenis karena melakukan analisis komparatif terhadap ketiga algoritma tersebut secara bersamaan dengan metode ADASYN pada konteks ulasan aplikasi *Access by KAI*, yang belum pernah dikaji dalam studi sebelumnya. Hasil penelitian menunjukkan bahwa setelah penerapan ADASYN, akurasi model menjadi 75,17% (SVM), 84,06% (RF), dan 83,17% (XGBoost). Meskipun akurasi sedikit menurun dibandingkan sebelum oversampling, nilai F1-score untuk kelas netral meningkat menjadi 0,13 (SVM), 0,05 (RF), dan 0,14 (XGBoost). Sebelumnya, tanpa ADASYN, akurasi mencapai 85,88% (SVM), 85,13% (RF), dan 85,37% (XGBoost), tetapi ketiganya gagal mengenali sentimen netral secara optimal dengan F1-score 0,00 (SVM dan RF) serta 0,03 (XGBoost). Hasil ini membuktikan bahwa penerapan ADASYN meningkatkan sensitivitas model terhadap sentimen netral, dengan XGBoost menjadi algoritma paling stabil dan optimal pada tugas klasifikasi sentimen aplikasi *Access by KAI*.

**Kata Kunci:** Access by KAI; Adaptive Synthetic Sampling; Support Vector Machine; Random Forest; XGBoost;

## Abstract

The rapid growth of digital applications has heightened the need to understand user perceptions more thoroughly, particularly through sentiment analysis of user-generated reviews. In practice, sentiment classification often faces challenges related to class imbalance, especially when neutral reviews are significantly fewer than positive or negative ones. This imbalance can limit a model's ability to accurately detect all sentiment categories. This study examines the comparative performance of three machine learning algorithms Support Vector Machine (SVM), Random Forest (RF), and Extreme Gradient Boosting (XGBoost) by applying the Adaptive Synthetic Sampling (ADASYN) technique to address class imbalance. This study differs from previous similar research by conducting a simultaneous comparative analysis of three algorithms using the ADASYN method in the context of Access by KAI application reviews, which has not been examined in prior studies. Experimental results indicate that after implementing ADASYN, model accuracies reached 75.17% for SVM, 84.06% for RF, and 83.17% for XGBoost. Although accuracy slightly decreased after oversampling, the F1-scores for the neutral class improved to 0.13 (SVM), 0.05 (RF), and 0.14 (XGBoost). Before applying ADASYN, the models achieved accuracies of 85.88% (SVM), 85.13% (RF), and 85.37% (XGBoost), but they were unable to effectively recognize neutral sentiments, with F1-scores of 0.00 for SVM and RF, and 0.03 for XGBoost. These findings suggest that ADASYN enhances model sensitivity to neutral sentiment, with XGBoost demonstrating the most consistent and robust performance in sentiment classification for the Access by KAI application.

**Keywords:** Access by KAI; Adaptive Synthetic Sampling; Support Vector Machine; Random Forest; XGBoost;

## 1. PENDAHULUAN

Penggunaan aplikasi *mobile* sebagai media utama pelayanan publik dalam beberapa tahun terakhir menunjukkan lonjakan perkembangan yang relatif pesat. Aplikasi *Access by KAI*, yang dimiliki oleh PT Kereta Api Indonesia, merupakan platform digital resmi yang menyediakan berbagai layanan seperti pemesanan tiket, pengecekan jadwal keberangkatan, hingga informasi perjalanan. Peningkatan pemanfaatan *Access by KAI* sejalan dengan tren digitalisasi pelayanan publik dan tingginya kebutuhan mobilitas masyarakat. Namun, masih ditemukan berbagai keluhan dari pengguna terkait performa sistem, seperti kesulitan login, gangguan pada transaksi pembelian tiket, hingga error pada jam sibuk [1]. Ulasan pengguna di Google Play Store menunjukkan variasi persepsi terhadap layanan ini, mulai dari yang sangat memuaskan hingga mengecewakan. Kondisi tersebut mengindikasikan perlunya analisis sentimen guna mendapatkan wawasan yang lebih komprehensif mengenai pengalaman pengguna dan tingkat kepuasan pengguna terhadap aplikasi *Access by KAI*.

Dalam bidang *machine learning*, analisis sentimen merupakan bagian dari *Natural Language Processing* (NLP) dan text mining yang berfokus pada pengenalan serta pengelompokan opini atau sikap seseorang terhadap objek, layanan, atau isu tertentu [2]. Sejumlah penelitian sebelumnya telah mengkaji sentimen pada ulasan pengguna *Access by KAI* dengan pendekatan yang bervariasi. Pada penelitian [3] digunakan metode *K-Nearest Neighbor* (KNN) dengan TF-IDF, namun belum mengatasi ketidakseimbangan data antara ulasan negatif (1.300) dan positif (457). Pada penelitian [4] dievaluasi algoritma *Logistic Regression*, *Naïve Bayes*, dan *Random Forest*, dengan akurasi terbaik pada *Logistic Regression* sebesar 84%, tetapi tanpa strategi penyeimbangan data. Penelitian [5] mengombinasikan teknik *word*

*embedding Word2Vec* dan SVM dengan kernel RBF, serta menerapkan SMOTE untuk menyeimbangkan data. Kombinasi ini menghasilkan akurasi sebesar 81% dan AUC 0,81, namun belum menguji ADASYN dan tidak membandingkannya dengan model ensemble seperti XGBoost. Penelitian [6] membandingkan performa dua kernel SVM (linear dan RBF) terhadap data ulasan pengguna *Access by KAI*, dengan akurasi terbaik 83,1% pada kernel linear, tetapi terbatas pada satu algoritma dan tidak menggunakan teknik *oversampling*. Penelitian [7] mengusulkan pendekatan berbasis leksikon dengan akurasi 87,2%, namun tanpa integrasi pembelajaran mesin serta penanganan ketidakseimbangan data. Sementara itu, pada konteks aplikasi lain seperti *Bumble*, penelitian [8] menunjukkan bahwa *Random Forest* unggul dibandingkan SVM dan XGBoost, tetapi belum mengimplementasikan metode penyeimbangan data seperti SMOTE atau ADASYN.

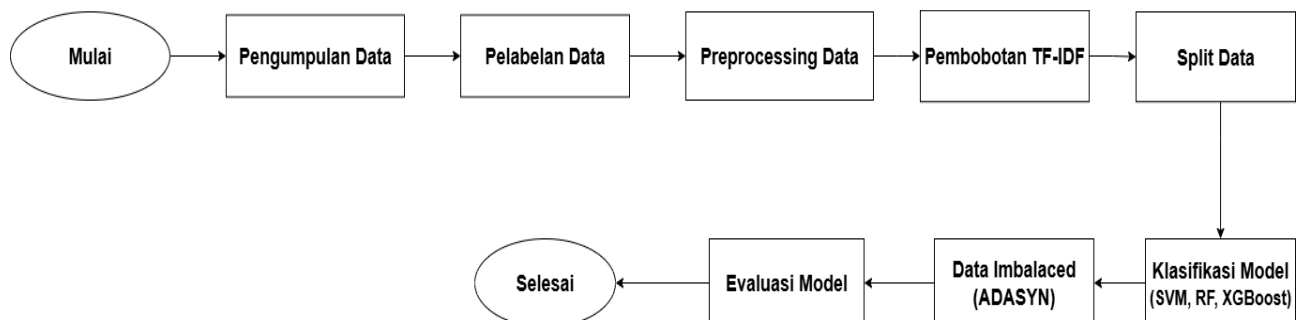
Ketidakseimbangan data merupakan tantangan umum dalam klasifikasi sentimen, terutama pada kelas netral yang sering kali terabaikan. Kondisi ini dapat menyebabkan model klasifikasi lebih condong ke kelas mayoritas, sehingga menurunkan kemampuan model dalam mendeteksi kelas minoritas. Pada penelitian [9] ditegaskan bahwa distribusi kelas yang tidak seimbang berdampak signifikan pada penurunan kinerja klasifikasi, khususnya terhadap data minoritas. Permasalahan ini dapat diatasi melalui penerapan metode *Adaptive Synthetic Sampling* (ADASYN), yaitu sebuah teknik yang secara adaptif menghasilkan data sintesis dengan mempertimbangkan tingkat kesulitan dalam mengklasifikasikan sampel dari kelas minoritas. ADASYN terbukti meningkatkan akurasi model klasifikasi pada data tidak seimbang. Pada penelitian [10] ditunjukkan bahwa ADASYN lebih efektif dibandingkan SMOTE dalam klasifikasi data kemiskinan. Penelitian [11] juga berhasil meningkatkan akurasi klasifikasi sentimen kenaikan harga BBM hingga 95% menggunakan ADASYN. Penelitian [12] menambahkan bahwa ADASYN memberikan performa stabil dan optimal pada rasio ketimpangan tinggi (rasio > 3), terutama saat digunakan bersama algoritma seperti *Random Forest* dan XGBoost.

Dari uraian diatas penelitian terdahulu yang mengkaji performa berbagai algoritma klasifikasi seperti SVM, Random Forest, dan XGBoost [3]–[8], serta studi mengenai efektivitas metode ADASYN dalam menangani ketidakseimbangan data [9]–[12], dapat disimpulkan bahwa penelitian sebelumnya belum pernah secara bersamaan membandingkan ketiga algoritma tersebut dengan penerapan ADASYN pada konteks ulasan aplikasi *Access by KAI*. Tiga algoritma yang digunakan dalam penelitian ini dipilih karena masing-masing memiliki keunggulan berbeda dalam menangani karakteristik data teks. SVM dikenal efektif dalam memisahkan data non-linear melalui pembentukan *hyperplane* optimal, RF juga dikenal unggul dalam mengurangi *overfitting* dengan pendekatan *ensemble* berbasis voting, sedangkan XGBoost dikenal mampu mengoptimalkan pembelajaran berurutan dengan akurasi tinggi pada dataset besar dan kompleks. Kombinasi ketiganya memberikan dasar yang kuat untuk analisis komparatif terhadap kinerja klasifikasi sentimen dengan data tidak seimbang. Kebaruan (*novelty*) penelitian ini terletak pada perbandingan kinerja SVM, RF, dan XGBoost secara bersamaan dengan penerapan teknik ADASYN pada konteks ulasan pengguna aplikasi *Access by KAI*. Belum ada penelitian sebelumnya yang secara spesifik membahas kombinasi ketiga algoritma ini menggunakan ADASYN untuk meningkatkan sensitivitas model terhadap kelas netral.

Fokus penelitian ini adalah membandingkan hasil perfoma tiga algoritma klasifikasi SVM, *Random Forest*, dan XGBoost dalam menentukan kategori sentimen terhadap ulasan pengguna *Access by KAI* untuk menilai algoritma mana yang paling stabil dan optimal. Teknik ADASYN digunakan untuk mengatasi ketidakseimbangan distribusi data, terutama pada kelas netral yang sering sulit dikenali oleh model. Sentimen netral penting untuk diidentifikasi karena mencerminkan opini moderat yang berisi masukan objektif, kritik membangun, atau pengalaman faktual yang bermanfaat bagi peningkatan kualitas layanan. Jika kelas ini terabaikan, analisis sentimen dapat menjadi bias terhadap opini ekstrem (positif dan negatif), sehingga mengurangi akurasi pemahaman terhadap persepsi pengguna secara keseluruhan. Penelitian ini memperluas penerapan ADASYN dalam analisis sentimen multi-kelas dan menilai pengaruhnya terhadap kinerja algoritma *machine learning* pada dataset yang tidak seimbang. Temuan studi ini diharapkan dapat menjadi acuan empiris bagi pengembang *Access by KAI* untuk memahami persepsi pengguna secara objektif serta merancang layanan digital dan sistem rekomendasi yang lebih responsif terhadap kebutuhan pengguna.

## 2. METODOLOGI PENELITIAN

Penelitian ini berfokus dengan membandingkan hasil kerja SVM, *Random Forest*, dan XGBoost dalam analisis sentimen ulasan pengguna aplikasi *Access by KAI*. Berikut skema penelitian secara keseluruhan diuraikan melalui Gambar 1.



**Gambar 1.** Skema Penelitian

## 2.1 Pengumpulan Data

Data diperoleh dengan cara mengimplementasikan teknik *web scraping* pada *Google Play Store* melalui *library google\_play\_scraper* [13].

## 2.2 Pelabelan Data

Data diberi label secara otomatis berdasarkan evaluasi pengguna, kemudian diklasifikasikan menjadi tiga tipe kelas sentimen: negatif, netral, dan positif. Ulasan dengan rating 1 dan 2 dimasukkan ke dalam kelas negatif karena mencerminkan tingkat kepuasan yang rendah. Rating 3 dikategorikan sebagai netral karena berada di tengah skala penilaian. Sementara itu, rating 4 dan 5 dianggap sebagai sentimen positif karena menunjukkan kepuasan pengguna. Metode pelabelan berdasarkan rating ini memberikan pendekatan yang objektif dalam menentukan kecenderungan sentimen berdasarkan penilaian pengguna secara langsung [14].

## 2.3 Preprocessing Data

Tahapan ini bertujuan untuk membersihkan dan menormalisasi teks sebelum dilakukan representasi fitur. Langkah-langkah preprocessing meliputi [15]:

- Cleansing* adalah tahapan pembersihan data, menghapus data duplikat, nilai kosong, simbol, angka, serta baris yang tidak relevan.
- Case Folding* adalah tahapan menstandarkan teks dengan mengkonversi semua karakter menjadi huruf kecil agar data tetap konsisten.
- Normalization* adalah tahapan untuk memperbaiki istilah atau teks yang tidak sesuai kaidah menjadi bentuk yang selaras dengan penulisan resmi.
- Tokenizing* adalah tahapan memecah kalimat atau kata menjadi unit-unit terkecil yang dikenal sebagai *token*.
- Stopword Removal* adalah tahapan penyaringan kata-kata *non-informatif* yang tidak menambah makna terhadap kalimat.
- Stemming* adalah tahapan yang memproses kata agar kembali ke bentuk dasar dengan menghilangkan awalan dan akhiran.

## 2.4 Pembobotan TF-IDF

Hasil prapemrosesan kemudian direpresentasikan secara numerik dengan penerapan metode (TF-IDF) *Term Frequency – Inverse Document Frequency*. Teknik ini merupakan pendekatan statistik untuk menilai seberapa penting suatu kata dalam sebuah dokumen relatif terhadap seluruh korpus. Komponen TF menunjukkan frekuensi kemunculan kata yang terdapat pada dokumen, sedangkan IDF menyeimbangkan bobot kata yang kemunculannya dominan di banyak dokumen [16].

## 2.5 Split Data

Setelah tahap ekstraksi fitur, *Split data* dilakukan ke dalam dua kelompok: *training set* dan *testing set*. Pemisahan ini berfungsi untuk membangun model dari sebagian data dan menguji kinerjanya pada data yang tidak dikenali oleh model. Strategi ini dikenal sebagai *train-test split* dan umumnya menggunakan rasio 80:20. Teknik ini juga diterapkan dalam penelitian [17] dalam set pelatihan dan pengujian guna melanjutkan proses klasifikasi analisis sentimen secara optimal.

## 2.6 Klasifikasi Model

Penilaian opini teks (*sentiment analysis*) adalah bagian dari NLP yang berfungsi untuk mendeteksi, menilai, dan mengklasifikasikan ekspresi emosi atau pendapat dalam teks menjadi kategori seperti negatif, netral, dan positif. Metode ini berada dalam domain komputasi afektif yang menekankan evaluasi informasi subjektif dari bahasa tulis [18]. Klasifikasi tersebut memerlukan algoritma yang mampu mengolah teks dengan baik agar kecenderungan opini dapat dianalisis secara sistematis. Tahap ini diarahkan untuk memilih algoritma klasifikasi paling optimal untuk ulasan pengguna aplikasi *Access by KAI*. Berikut penjelasan singkat dari masing-masing algoritma:

- Support Vector Machine (SVM)* adalah algoritma yang membagi data melalui pembentukan *hyperplane* optimal untuk membedakan kelas-kelas berbeda dalam ruang berdimensi tinggi. Algoritma ini efektif untuk menangani data besar dan teks tidak terstruktur. *Hyperplane* yang dibentuk SVM memaksimalkan jarak antar kelas (*margin*) sehingga model memiliki kemampuan generalisasi yang tinggi terhadap data baru. Objek-objek data yang terletak paling berdekatan dengan *hyperplane* disebut *support vectors*, yang berperan penting dalam penentuan posisi *hyperplane* [19]. Pada studi ini, SVM dikonfigurasi dengan *kernel*='linear' dan *random\_state*=42. Kernel linear dipilih karena representasi fitur TF-IDF menghasilkan ruang berdimensi tinggi, sehingga linear SVM cukup efektif sekaligus efisien secara komputasi.
- Random Forest* adalah metode *ensemble learning* yang memanfaatkan banyak *decision tree*. Setiap pohon dibangun dari sampel acak data pelatihan, dan hasil prediksi akhir ditentukan melalui *majority voting*. *Random Forest* efektif untuk menangani masalah klasifikasi maupun regresi, serta bekerja baik pada dataset yang tidak seimbang (*imbalanced data*) [20]. Model RF pada studi ini menggunakan *n\_estimators*=100, *criterion*='gini', dan *random\_state*=42, yang memberikan keseimbangan antara akurasi dan waktu pelatihan.

- c. *Extreme Gradient Boosting* (XGBoost) adalah algoritma gradient boosting yang membangun model secara berurutan, dengan setiap model baru fokus memperbaiki kesalahan prediksi model sebelumnya. Algoritma ini efisien dalam hal kecepatan dan penggunaan memori, serta mampu memberikan performa klasifikasi yang stabil dan akurat, bahkan tanpa teknik penyeimbangan data [21]. Dalam studi ini, XGBoost diatur dengan `use_label_encoder=False`, `eval_metric='mlogloss'`, `random_state=42`, dan `n_estimators=100`, untuk menjaga stabilitas model dan mengoptimalkan proses pelatihan.

## 2.7 Teknik Oversampling (ADASYN)

Metode *oversampling* ADASYN (*Adaptive Synthetic Sampling*) diterapkan guna menangani distribusi kelas yang tidak seimbang pada dataset. ADASYN diimplementasikan guna menghasilkan data sintesis dari kelas minoritas yang sulit diprediksi oleh model, berdasarkan distribusi lokalnya. Pada penelitian [22] membuktikan bahwa penerapan ADASYN mampu meningkatkan akurasi model *Hard Voting Classifier* dalam klasifikasi penyakit diabetes dari 99% menjadi 99,8%. Selain itu, dalam penelitian [23] pada klasifikasi komentar tentang program Kampus Merdeka menunjukkan bahwa kombinasi ADASYN dengan TF-IDF dan SVM dapat meningkatkan performa klasifikasi secara signifikan, khususnya pada nilai F1-score. Untuk memperjelas dampak penerapan metode ADASYN, penelitian ini juga membandingkan distribusi jumlah data antar kelas sebelum dan sesudah proses *oversampling*. Setelah diterapkan, distribusi data di setiap kelas menjadi lebih merata, memungkinkan model untuk lebih efektif mengidentifikasi pola sentimen netral.

## 2.8 Evaluasi Model

Setelah model klasifikasi dibangun, langkah berikutnya adalah mengevaluasi performanya. Evaluasi dilakukan menggunakan data uji yang diperoleh dari proses *train-test split*, untuk mengukur akurasi model dalam mengklasifikasikan data baru yang sebelumnya belum diproses, tahap evaluasi dalam sistem analisis sentimen menjadi aspek penting untuk menilai keandalan klasifikasi opini. Proses penilaian dilakukan melalui metrik presisi, recall, F1-score, dan akurasi, keempat metrik tersebut dihitung berdasarkan nilai *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) yang diperoleh melalui *confusion matrix* [24]. Adapun rumus yang digunakan adalah sebagai berikut:

$$\text{Precision} = \frac{TP}{TP+FP} \quad (1)$$

$$\text{Recall} = \frac{TP}{TP+FN} \quad (2)$$

$$\text{F1-Score} = \frac{2 \times \text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}} \quad (3)$$

$$\text{Accuracy} = \frac{TP_A + TP_B + TP_C}{TP_A + TP_B + TP_C + FP_A + FP_B + FP_C + FN_A + FN_B + FN_C} \quad (4)$$

# 3. HASIL DAN PEMBAHASAN

Bagian ini memaparkan hasil pengujian serta penilaian hasil kerja tiga algoritma *machine learning*, yaitu —SVM, *Random Forest*, dan XGBoost—yang diterapkan untuk menganalisis sentimen pada data ulasan aplikasi *Access by KAI* di Google Play Store. Fokus utama terletak pada perbandingan kinerja ketiganya sebelum dan sesudah penerapan teknik *oversampling* ADASYN, yang bertujuan mengatasi ketidakseimbangan distribusi data antar kelas sentimen minoritas. Selain itu, ditampilkan pula perbandingan distribusi label sebelum dan sesudah ADASYN untuk mengilustrasikan perubahan proporsi kelas karena penerapan teknik ADASYN.

## 3.1 Pengumpulan Data

Dari hasil scraping secara keseluruhan, terkumpul sebanyak 108.904 data ulasan berupa opini pengguna aplikasi *Access by KAI*. Namun, untuk menjaga relevansi temporal dan kualitas analisis, hanya ulasan yang ditulis pada tahun 2025 yang disertakan dalam penelitian ini. Setelah dilakukan penyaringan berdasarkan tahun publikasi, diperoleh sebanyak 10.693 ulasan yang digunakan sebagai dataset utama. Distribusi awal data tersebut terdiri dari 5.234 ulasan dengan sentimen negatif, 4.885 sentimen positif, dan hanya 574 sentimen netral. Komposisi ini menunjukkan adanya ketidakseimbangan kelas (*imbalanced class distribution*) yang cukup signifikan, terutama pada kelas netral yang berjumlah jauh lebih sedikit dibanding kelas lainnya.

## 3.2 Preprocessing

Setelah proses pengumpulan data, langkah berikutnya adalah *preprocessing* guna menyiapkan teks agar siap dianalisis lebih lanjut. Data yang dihasilkan pada tahap ini dipaparkan pada Tabel 1.

**Tabel 1.** Hasil *Preprocessing*

| Proses               | Hasil                               |
|----------------------|-------------------------------------|
| <i>Cleansing</i>     | Praktis dan hemat biaya plus energi |
| <i>Case Folding</i>  | praktis dan hemat biaya plus energi |
| <i>Normalization</i> | praktis dan hemat biaya plus energi |



|                  |                                            |
|------------------|--------------------------------------------|
| Tokenizing       | [praktis, dan, hemat, biaya, plus, energi] |
| Stopword Removal | [praktis, hemat, biaya, plus, energi]      |
| Steming          | praktis hemat biaya plus energi            |

Berdasarkan Tabel 1, hasil *preprocessing* menunjukkan bahwa teks telah dibersihkan dan distandarisasi sehingga lebih terstruktur dan setiap kalimat dapat direpresentasikan sebagai token yang bermakna. Proses ini berhasil menghapus kata-kata yang tidak relevan, menyamakan format huruf, serta menyederhanakan kata ke bentuk dasarnya, sehingga mempermudah model dalam mempelajari pola kata dan mendukung akurasi pada klasifikasi sentimen.

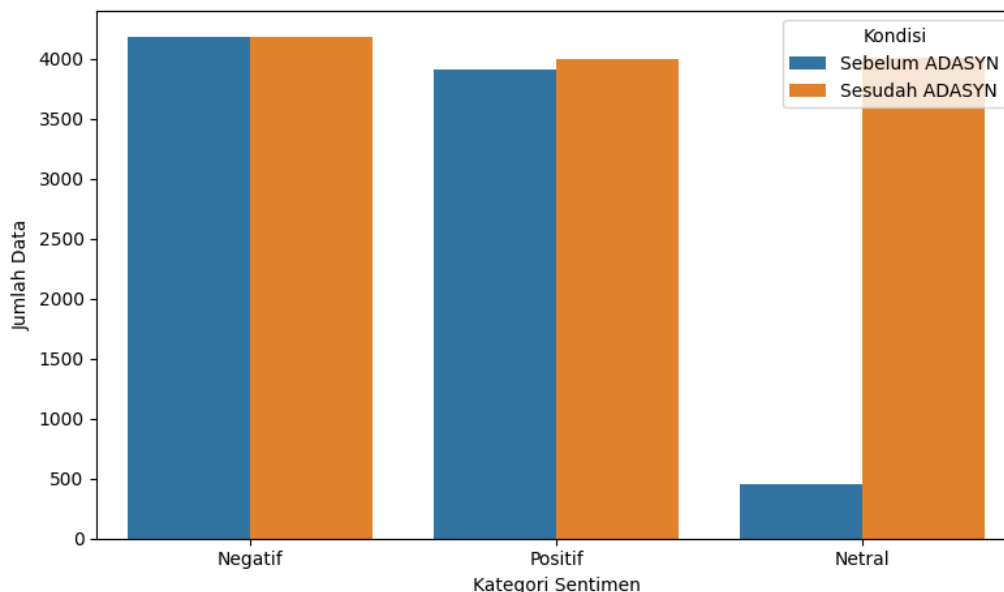
### 3.4 Split Data

Dataset dipisahkan menjadi dua kelompok menggunakan pendekatan *stratified split*, dengan 80% digunakan sebagai data untuk pelatihan 8.554 sampel dan 20% sebagai data untuk pengujian 2.139 sampel.

### 3.5 ADASYN

Distribusi data hasil *preprocessing* dan pelabelan menunjukkan ketidakseimbangan kelas yang cukup signifikan, dengan kelas netral memiliki jumlah yang jauh lebih sedikit dibandingkan kelas negatif maupun positif. Sebelum dilakukan *oversampling*, data latih terdiri atas 4.187 data negatif, 3.908 positif, dan hanya 459 netral. Ketimpangan ini berpotensi menyebabkan model lebih cenderung mengklasifikasikan ulasan ke dalam kelas mayoritas dan mengabaikan kelas minoritas, sehingga menurunkan performa model secara keseluruhan.

Sebagai upaya mengatasi masalah tersebut, penelitian ini menerapkan pendekatan ADASYN, yaitu teknik *oversampling* yang menghasilkan data sintesis dari kelas minoritas berdasarkan tingkat kesulitannya untuk diklasifikasikan [10]. Metode ini bertujuan untuk memperkuat representasi kelas minoritas, khususnya sentimen netral, agar model memiliki kemampuan klasifikasi yang lebih seimbang. Berikut visualisasi perbandingan distribusi data sebelum dan sesudah penerapan ADASYN untuk menilai efek *oversampling* terhadap ketidakseimbangan kelas yang ditampilkan pada Gambar 2.



**Gambar 2.** Visualisasi Perbandingan Distribusi Data Sebelum dan Sesudah ADASYN

Visualisasi pada Gambar 2 memperlihatkan perubahan proporsi antar kelas sentimen setelah proses *oversampling*, di mana kelas netral menjadi lebih seimbang terhadap kelas negatif dan positif. Setelah penerapan ADASYN, distribusi data menjadi lebih proporsional: 4.187 data negatif, 4.007 netral, dan 3.995 positif. Hasil penelitian ini sejalan dengan temuan sebelumnya, yang menegaskan bahwa ADASYN mampu meningkatkan kinerja model klasifikasi, terutama pada algoritma pohon keputusan seperti *Random Forest* dan *XGBoost* [22]. Dengan demikian, teknik *oversampling* ini terbukti efektif memperkuat representasi kelas minoritas, sehingga model memiliki performa klasifikasi yang lebih proporsional dan adil terhadap seluruh kategori sentimen.

### 3.6 Evaluasi Model

Bagian ini menyajikan hasil evaluasi performa dari tiga algoritma klasifikasi, yakni SVM, Random Forest, dan XGBoost, yang diaplikasikan pada penelitian kali ini. Untuk mengevaluasi efektivitas model pada tugas klasifikasi sentimen, digunakan indikator akurasi, presisi, recall, dan F1-score. Proses evaluasi membandingkan dua skenario, yakni sebelum dan setelah penerapan teknik *oversampling* ADASYN guna mengatasi ketidakseimbangan distribusi data antar kelas.

### 3.6.1 Evaluasi Model Sebelum Penerapan Teknik *Oversampling* ADASYN

Pada bagian ini membahas performa awal dari masing-masing model klasifikasi ketika diterapkan pada data ulasan pengguna sebelum penerapan teknik *oversampling* ADASYN. Evaluasi dilakukan untuk menilai kinerja model dalam mendeteksi sentimen pada data yang memiliki distribusi kelas tidak seimbang, terutama pada kelas minoritas seperti sentimen netral. Detail evaluasi tiap algoritma sebelum penerapan ADASYN dipaparkan pada Tabel 2, 3, dan 4.

**Tabel 2.** Hasil Evaluasi Performa Algoritma SVM Sebelum Penerapan Teknik *Oversampling* ADASYN

|              | Precision | Recall | F-1 Score | Support |
|--------------|-----------|--------|-----------|---------|
| Negatif      | 0.82      | 0.94   | 0.88      | 1047    |
| Netral       | 0.00      | 0.00   | 0.00      | 115     |
| Positif      | 0.91      | 0.88   | 0.89      | 977     |
| Accuracy     |           |        | 0.86      | 2139    |
| Macro Avg    | 0.58      | 0.60   | 0.59      | 2139    |
| Weughted Avg | 0.82      | 0.86   | 0.84      | 2139    |

Berdasarkan Tabel 2, algoritma SVM menunjukkan performa tinggi pada kelas negatif dan positif, dengan precision masing-masing 0,82 dan 0,91. Namun, performa pada kelas netral sangat rendah (*precision, recall, F1-score* = 0), mengindikasikan bahwa SVM kesulitan mengenali sentimen minoritas sebelum penerapan ADASYN. Secara keseluruhan, akurasi SVM pada data sebelum *oversampling* adalah 85,88%.

**Tabel 3.** Hasil Evaluasi Performa Algoritma *Random Forest* Sebelum Penerapan Teknik *Oversampling* ADASYN

|              | Precision | Recall | F-1 Score | Support |
|--------------|-----------|--------|-----------|---------|
| Negatif      | 0.79      | 0.96   | 0.87      | 1047    |
| Netral       | 0.00      | 0.00   | 0.00      | 115     |
| Positif      | 0.94      | 0.84   | 0.89      | 977     |
| Accuracy     |           |        | 0.85      | 2139    |
| Macro Avg    | 0.58      | 0.60   | 0.58      | 2139    |
| Weughted Avg | 0.82      | 0.85   | 0.83      | 2139    |

Tabel 3 memperlihatkan pola serupa pada *Random Forest*. Kelas negatif dan positif memiliki performa cukup baik, tetapi kelas netral tetap tidak terdeteksi (nilai 0), menegaskan adanya bias terhadap kelas mayoritas sebelum *oversampling*. Akurasi keseluruhan *Random Forest* tercatat 85,13%.

**Tabel 4.** Hasil Evaluasi Performa Algoritma XG-Boost Sebelum Penerapan Teknik *Oversampling* ADASYN

|              | Precision | Recall | F-1 Score | Support |
|--------------|-----------|--------|-----------|---------|
| Negatif      | 0.82      | 0.93   | 0.87      | 1047    |
| Netral       | 0.33      | 0.02   | 0.03      | 115     |
| Positif      | 0.90      | 0.88   | 0.89      | 977     |
| Accuracy     |           |        | 0.85      | 2139    |
| Macro Avg    | 0.68      | 0.61   | 0.60      | 2139    |
| Weughted Avg | 0.83      | 0.85   | 0.83      | 2139    |

Tabel 4 menunjukkan bahwa XGBoost sedikit lebih mampu mengenali kelas netral dibanding SVM dan *Random Forest*, dengan *recall* 0,02 dan *precision* 0,33. Meskipun ada peningkatan, kelas minoritas masih terabaikan, sehingga diperlukan teknik *oversampling* seperti ADASYN untuk memperbaiki distribusi kelas dan meningkatkan performa model. Akurasi keseluruhan XGBoost pada data sebelum *oversampling* adalah 85,37%.

### 3.6.2 Evaluasi Model Setelah Penerapan Teknik *Oversampling* ADASYN

Setelah klasifikasi tanpa *oversampling*, hasil evaluasi menunjukkan ketiga algoritma mengalami kesulitan signifikan dalam mendeteksi sentimen netral. Pada bagian ini dibahas performa masing-masing model klasifikasi setelah penerapan teknik *oversampling* ADASYN. Evaluasi dilakukan untuk menilai efektivitas model dalam mengenali sentimen pada data ulasan pengguna setelah distribusi kelas diseimbangkan, khususnya pada kelas minoritas seperti sentimen netral. Detail evaluasi tiap algoritma setelah penerapan ADASYN dipaparkan pada Tabel 5, 6, dan 7.

**Tabel 5.** Hasil Evaluasi Performa Algoritma SVM Setelah Penerapan Teknik *Oversampling* ADASYN

|           | Precision | Recall | F-1 Score | Support |
|-----------|-----------|--------|-----------|---------|
| Negatif   | 0.84      | 0.83   | 0.83      | 1047    |
| Netral    | 0.09      | 0.27   | 0.13      | 115     |
| Positif   | 0.94      | 0.73   | 0.82      | 977     |
| Accuracy  |           |        | 0.75      | 2139    |
| Macro Avg | 0.62      | 0.61   | 0.60      | 2139    |

|              |      |      |      |      |
|--------------|------|------|------|------|
| Weighted Avg | 0.84 | 0.75 | 0.79 | 2139 |
|--------------|------|------|------|------|

Berdasarkan Tabel 5, SVM menunjukkan peningkatan kemampuan mengenali kelas minoritas setelah penerapan ADASYN. F1-score kelas netral naik menjadi 0.13 (*precision* 0.09, *recall* 0.27), meskipun masih lebih rendah dibanding kelas mayoritas. Akurasi keseluruhan SVM menjadi 75,17%.

**Tabel 6.** Hasil Evaluasi Performa Algoritma *Random Forest* Setelah Penerapan Teknik *Oversampling* ADASYN

|              | Precision | Recall | F-1 Score | Support |
|--------------|-----------|--------|-----------|---------|
| Negatif      | 0.80      | 0.93   | 0.86      | 1047    |
| Netral       | 0.10      | 0.03   | 0.05      | 115     |
| Positif      | 0.92      | 0.83   | 0.87      | 977     |
| Accuracy     |           |        | 0.84      | 2139    |
| Macro Avg    | 0.61      | 0.60   | 0.59      | 2139    |
| Weighted Avg | 0.82      | 0.84   | 0.82      | 2139    |

Tabel 6 menunjukkan *Random Forest* mengalami peningkatan F1-score kelas netral menjadi 0.05 (*precision* 0.10, *recall* 0.03). Nilai ini masih rendah, tetapi representasi kelas minoritas lebih baik dibanding sebelum *oversampling*. Akurasi keseluruhan meningkat menjadi 84,06%.

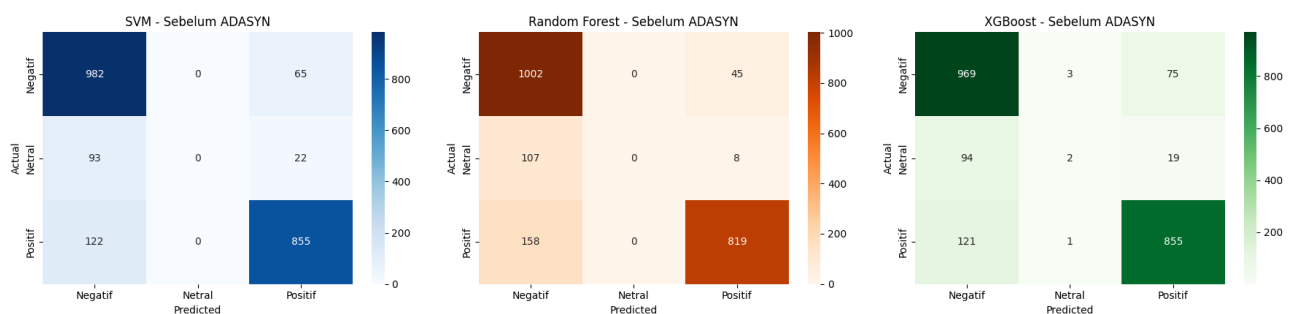
**Tabel 7.** Hasil Evaluasi Performa Algoritma XG-Boost Setelah Penerapan Teknik *Oversampling* ADASYN

|              | Precision | Recall | F-1 Score | Support |
|--------------|-----------|--------|-----------|---------|
| Negatif      | 0.82      | 0.87   | 0.84      | 1047    |
| Netral       | 0.15      | 0.13   | 0.14      | 115     |
| Positif      | 0.9       | 0.87   | 0.89      | 977     |
| Accuracy     |           |        | 0.83      | 2139    |
| Macro Avg    | 0.63      | 0.62   | 0.62      | 2139    |
| Weighted Avg | 0.83      | 0.83   | 0.83      | 2139    |

Tabel 7 memperlihatkan XGBoost mengalami peningkatan performa pada kelas netral menjadi *F1-score* 0.14 (*precision* 0.15, *recall* 0.13). Akurasi keseluruhan XGBoost setelah ADASYN tercatat 83,17%, menunjukkan distribusi kelas yang lebih seimbang meningkatkan kemampuan model mengenali seluruh kategori sentimen.

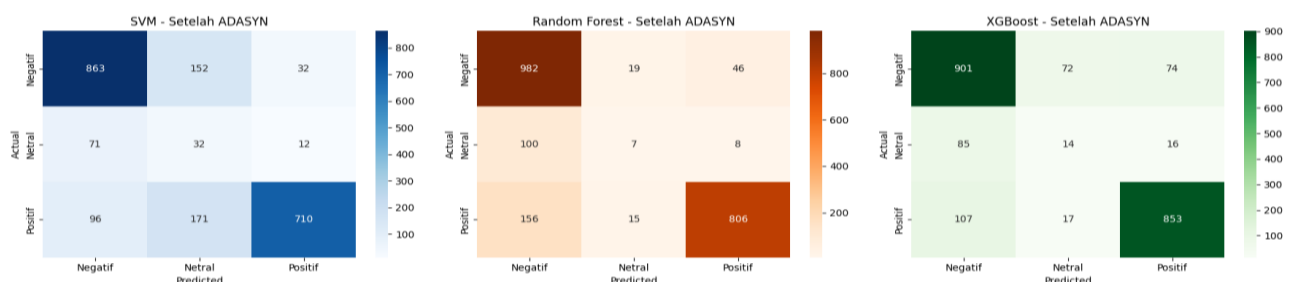
### 3.6.3 Confusion Matrix

Visualisasi *confusion matrix* diterapkan untuk membandingkan kinerja model SVM, *Random Forest*, dan XGBoost sebelum dan sesudah penerapan teknik *oversampling* ADASYN, ditampilkan pada Gambar 3 dan Gambar 4.



**Gambar 3.** Confusion Matrix Sebelum Penerapan Teknik *Oversampling* ADASYN

Berdasarkan Gambar 3, sebelum penerapan ADASYN, ketiga model cenderung bias terhadap kelas mayoritas (Negatif dan Positif) dengan tingkat prediksi yang sangat rendah pada kelas netral. Pada SVM, tidak ada satu pun data kelas netral yang berhasil diklasifikasikan dengan benar (prediksi = 0). Hal serupa juga terlihat pada *Random Forest* dan XGBoost, yang mengindikasikan bahwa model-model ini memiliki keterbatasan dalam menangani kelas minoritas.



**Gambar 4.** Confusion Matrix Sesudah Penerapan Teknik *Oversampling* ADASYN

Setelah penerapan ADASYN, terlihat pada Gambar 4 bahwa distribusi prediksi antar kelas menjadi lebih seimbang. SVM mulai mampu memprediksi sebagian data kelas netral, *Random Forest* meningkatkan akurasi pada kelas netral sekaligus mengurangi kesalahan pada kelas lain, sementara XGBoost menunjukkan peningkatan paling signifikan. Secara umum, ADASYN terbukti memperkuat representasi kelas minoritas sehingga kinerja model menjadi lebih adil dan akurat, meskipun belum sepenuhnya sempurna. Berdasarkan *confusion matrix* diatas, dapat diamati perubahan distribusi prediksi model sebelum dan sesudah penerapan teknik oversampling ADASYN. Untuk memberikan gambaran yang lebih komprehensif mengenai perbandingan performa klasifikasi, khususnya dalam mendeteksi kelas sentimen netral, hasil evaluasinya dipaparkan secara rinci pada Tabel 8.

**Tabel 8.** Perbandingan Kinerja Klasifikasi Sentimen Netral Sebelum dan Sesudah Penerapan ADASYN

|                        | Precision | Recall | F-1 Score | Support |
|------------------------|-----------|--------|-----------|---------|
| SVM                    | 0.00      | 0.00   | 0.00      | 115     |
| SVM (ADASYN)           | 0.09      | 0.27   | 0.13      | 115     |
| Random Forest          | 0.00      | 0.00   | 0.00      | 115     |
| Random Forest (ADASYN) | 0.10      | 0.03   | 0.05      | 115     |
| XGBoost                | 0.33      | 0.02   | 0.03      | 115     |
| XGBoost (ADASYN)       | 0.15      | 0.13   | 0.14      | 115     |

Berdasarkan Tabel 8, dapat dilihat bahwa penerapan teknik oversampling ADASYN memberikan peningkatan nyata pada performa klasifikasi sentimen netral di ketiga algoritma. Nilai *F1-score* kelas netral yang semula sangat rendah mulai menunjukkan perbaikan, terutama pada SVM dan XGBoost, yang mengindikasikan bahwa ADASYN mampu memperkuat representasi data minoritas tanpa menurunkan akurasi secara signifikan. Berbeda dengan penelitian sebelumnya yang menganalisis ulasan aplikasi *Access by KAI* menggunakan satu algoritma tanpa melakukan perbandingan model, penelitian ini menerapkan pendekatan yang lebih komprehensif dengan membandingkan hasil kinerja tiga algoritma klasifikasi SVM, *Random Forest*, dan XGBoost yang dikombinasikan dengan penerapan teknik *oversampling* ADASYN. Penelitian sebelumnya menggunakan algoritma SVM dengan metode SMOTE untuk menangani ketidakseimbangan data, namun belum melakukan evaluasi dengan membandingkan lebih dari satu model dan belum menyoroti secara spesifik performa terhadap kelas sentimen netral yang merupakan kategori minoritas [5].

### 3.6.4 Visualisasi Word Cloud

Untuk memperkuat analisis klasifikasi, dilakukan visualisasi kata-kata dengan frekuensi tertinggi pada tiap kelas sentimen dengan teknik *Word Cloud*. Visualisasi ini berfungsi untuk mengungkap kata-kata yang menjadi fokus utama, serta mengidentifikasi keluhan dan pujian yang paling banyak diungkapkan oleh pengguna aplikasi *Access by KAI*. Berikut visualisasi *word cloud* yang menunjukkan kosakata paling dominan pada masing-masing kelas sentimen ditampilkan pada Gambar 5, Gambar 6, dan Gambar 7.



**Gambar 5.** Visualisasi *Word Cloud* Sentimen Negatif

Berdasarkan Gambar 5, *Word Cloud* untuk sentimen negatif menunjukkan kata yang paling dominan, seperti “error”, “gagal”, “tidak bisa”, dan “lagi”. Hal ini mengindikasikan masalah-masalah yang sering dialami pengguna, misalnya kegagalan transaksi atau kesulitan dalam menggunakan fitur aplikasi.



**Gambar 6.** Visualisasi *Word Cloud* Sentimen Netral



Pada Gambar 6, *Word Cloud* sentimen netral didominasi kata-kata seperti “saya”, “bisa”, dan “tiket”. Hal ini menunjukkan bahwa banyak pengguna menyampaikan informasi atau pengalaman yang bersifat deskriptif dan tidak mengekspresikan kepuasan maupun ketidakpuasan secara eksplisit.



**Gambar 7.** Visualisasi *Word Cloud* Sentimen Positif

Berdasarkan Gambar 7 menampilkan *Word Cloud* untuk sentimen positif, yang menonjolkan kata-kata seperti “bagus”, “mantap”, “cepat”, dan “terima kasih”. Kata-kata ini mencerminkan kepuasan pengguna terhadap kinerja aplikasi, kecepatan layanan, serta pengalaman penggunaan yang menyenangkan. Dengan demikian, visualisasi *Word Cloud* memberikan gambaran jelas mengenai kosakata dominan di tiap kelas sentimen, mendukung pemahaman lebih mendalam terhadap keluhan, informasi netral, dan pujian pengguna.

#### 4. KESIMPULAN

Hasil perbandingan kinerja tiga algoritma klasifikasi ini menunjukkan bahwa penerapan teknik ADASYN sangat efektif dalam mengatasi ketidakseimbangan data pada klasifikasi sentimen ulasan aplikasi *Access by KAI*. Setelah penerapan teknik ADASYN, terjadi peningkatan kemampuan ketiga algoritma klasifikasi dalam mengenali kelas netral pada ulasan aplikasi *Access by KAI*. F1-score meningkat menjadi 0,13 pada SVM, 0,05 pada RF, dan 0,14 pada XGBoost. Meskipun akurasi keseluruhan mengalami penurunan, hal ini merupakan *trade-off* yang dapat diterima karena model menjadi lebih seimbang dan mampu mendeteksi sentimen netral yang sebelumnya sulit dikenali secara akurat dan konsisten. Sebelum penerapan ADASYN, ketiga metode memiliki akurasi keseluruhan yang cukup tinggi, yaitu 85,88% pada SVM, 85,13% pada RF, dan 85,37% pada XGBoost, namun gagal mengenali kelas netral secara optimal. Hal ini ditandai dengan F1-score sebesar 0,00 untuk SVM dan RF, serta 0,03 pada XGBoost, yang mengindikasikan bahwa model-model tersebut bias terhadap kelas mayoritas dan kurang sensitif terhadap kelas minoritas, sehingga performa dalam mendeteksi kelas netral sangat terbatas. Dari perbandingan ketiga algoritma ini, XGBoost merupakan metode yang paling optimal dan stabil pada ulasan aplikasi *Access by KAI*. Setelah ADASYN diterapkan, XGBoost tidak hanya mempertahankan akurasi yang tinggi, tetapi juga memperlihatkan peningkatan paling signifikan dalam kemampuan mengenali kelas minoritas secara efektif. Hal ini menjadikan XGBoost pilihan paling optimal dan andal dalam menghadapi masalah ketidakseimbangan data pada tugas klasifikasi sentimen ini. Penelitian ini menunjukkan bahwa penerapan teknik *oversampling* ADASYN dapat meningkatkan sensitivitas model terhadap kelas minoritas dalam tugas klasifikasi sentimen. Temuan ini dapat menjadi pedoman bagi pengembang aplikasi dan analis data dalam memilih model yang lebih seimbang dan andal untuk memahami sentimen pengguna. Pada studi berikutnya, disarankan untuk memperluas eksperimen dengan algoritma yang berbeda atau pendekatan pembelajaran mendalam (*deep learning*) agar deteksi kelas minoritas lebih optimal dan generalisasi model terhadap berbagai dataset meningkat.

## REFERENCES

- [1] R. Damanhuri and V. A. Husein, "Analisis Sentimen pada Ulasan Aplikasi *Access by KAI* Berbahasa Indonesia Menggunakan Word-Embedding dan Classical Machine Learning," *J. Masy. Inform.*, vol. 15, no. 2, pp. 97–106, 2024, doi: 10.14710/jmasif.15.2.62383.
- [2] D. Purnamasari *et al.*, *Pengantar Metode Analisis Sentimen*. 2023.
- [3] N. A. Jiana and B. Hartono, "Sentimen Analisa Ulasan Aplikasi *Access by KAI* pada Google Play Store menggunakan Algoritma K-NN," *J. Media Inform. Budidarma*, vol. 8, no. 3, p. 1388, 2024, doi: 10.30865/mib.v8i3.7730.
- [4] M. A. S. Nugroho, D. Susilo, and D. Retnoningsih, "Analisis Sentimen Ulasan Aplikasi "Access by KAI" Menggunakan Algoritma Machine Learning," *J. Tek. Inf. dan Komput.*, vol. 7, no. 2, p. 820, 2024, doi: 10.37600/tekinkom.v7i2.1854.
- [5] D. L. Devi, A. A. Arifiyanti, and S. F. A. Wati, "Analisis Sentimen Ulasan Pengguna *Access by KAI* Menggunakan Metode Word2Vec Dan Algoritma Svm," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, 2024, doi: 10.23960/jitet.v12i3.4892.
- [6] E. Rizqi Mar'atus Sholihah, I. G. Susrama Mas Diyasa, and E. Yulia Puspaningrum, "Perbandingan Kinerja Kernel Linear Dan Rbf Support Vector Machine Untuk Analisis Sentimen Ulasan Pengguna KAI Access Pada Google Play Store," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 1, pp. 728–733, 2024, doi: 10.36040/jati.v8i1.8800.
- [7] R. D. Wahyuni and A. N. Utomo, "Penggunaan Metode Lexicon Untuk Analisis Sentimen pada Ulasan Aplikasi KAI Access di

- Google Play Store,” *J. Rekayasa Inf.*, vol. 11, no. 2, pp. 134–145, 2022.
- [8] O. Oktafia and R. S. A. Nugroho, “Comparison of Support Vector Machine(Svm), Xgboost and *Random Forest* for Sentiment Analysis of Bumble App User Comments,” *Proxies J. Inform.*, vol. 6, no. 1, pp. 32–46, 2024, doi: 10.24167/proxies.v6i1.12453.
- [9] I Gusti Ngurah Ady Kusuma, I Made Pradipta, I Made Ari Santosa, and I Komang Dharmendra, “Penanganan Ketidakseimbangan Data Pada Klasifikasi Pengaduan Masyarakat,” *J. Teknol. Inf. dan Komput.*, vol. 9, no. 5, pp. 489–496, 2023, doi: 10.36002/jutik.v9i5.2643.
- [10] D. V. Ramadhanti, R. Santoso, and T. Widiari, “Perbandingan Smote Dan Adasyn Pada Data Imbalance Untuk Klasifikasi Rumah Tangga Miskin Di Kabupaten Temanggung Dengan Algoritma K-Nearest Neighbor,” *J. Gaussian*, vol. 11, no. 4, pp. 499–505, 2023, doi: 10.14710/j.gauss.11.4.499-505.
- [11] F. Fauzi, I. Ismatullah, and I. M. Nur, “Adaptive Synthetic Support Vector Machine Multiclass untuk mengklasifikasikan Imbalance data pada Sentimen kenaikan Bahan Bakar Minyak,” *Pros. Semin. Nas. Sains Data*, vol. 3, no. 1, pp. 304–312, 2023, doi: 10.33005/senada.v3i1.127.
- [12] M. Imani, A. Beikmohammadi, and H. R. Arabnia, “Comprehensive Analysis of *Random Forest* and XGBoost Performance with SMOTE, ADASYN, and GNUS Under Varying Imbalance Levels,” *Technologies*, vol. 13, no. 3, pp. 1–40, 2025, doi: 10.3390/technologies13030088.
- [13] S. A. S. Mola, Y. C. Luttu, and D. N. Rumklakl, “Perbandingan Metode Machine Learning dalam Analisis Sentimen Komentar Pengguna Aplikasi InDriver pada Dataset Tidak Seimbang,” *J. Sist. Inf. Bisnis*, vol. 14, no. 3, pp. 247–255, 2024, doi: 10.21456/vol14iss3pp247-255.
- [14] H. N. Zuhdi and B. Prasetyo, “Analisis Sentimen pada Ulasan Aplikasi iPusnas di Google Play Store Menggunakan Naive Bayes Classifier,” *J. Homepage https://journal.irpi.or.id/index.php/ijirse*, vol. 5, no. 1, pp. 12–19, 2025.
- [15] S. Nachwa *et al.*, “Pendekatan Klasifikasi Dalam Knowledge Discovery Untuk Analisis Sentimen Berbasis Aspek Pada Ulasan Bandara Sultan Mahmud Badaruddin II Di Google Maps,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 3, pp. 4782–4789, 2025, doi: 10.36040/jati.v9i3.13776.
- [16] I. G. B. A. Budaya and I. K. P. Suniantara, “Comparison of Sentiment Analysis Algorithms with SMOTE Oversampling and TF-IDF Implementation on Google Reviews for Public Health Centers,” *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 3, pp. 1077–1086, 2024, doi: 10.57152/malcom.v4i3.1459.
- [17] E. D. Madyatmadja, Shinta, D. Susanti, F. Anggreani, and D. J. M. Sembiring, “Sentiment Analysis on User Reviews of Mutual Fund Applications,” *J. Comput. Sci.*, vol. 18, no. 10, pp. 885–895, 2022, doi: 10.3844/jcssp.2022.885.895.
- [18] M. Kumar, L. Khan, and H. T. Chang, “Evolving techniques in sentiment analysis: a comprehensive review,” *PeerJ Comput. Sci.*, vol. 11, pp. 1–61, 2025, doi: 10.7717/PEERJ-CS.2592.
- [19] D. A. Agustina, S. Subanti, and E. Zukhrinah, “Implementasi Text Mining Pada Analisis Sentimen Pengguna Twitter Terhadap Marketplace di Indonesia Menggunakan Algoritma Support Vector Machine,” *Indones. J. Appl. Stat.*, vol. 3, no. 2, p. 109, 2021, doi: 10.13057/ijas.v3i2.44337.
- [20] A. Nugroho and A. Husin, “Performance Analysis of *Random Forest* Using Attribute Normalization,” *Sistemasi*, vol. 11, no. 1, p. 186, 2022, doi: 10.32520/stmsi.v11i1.1681.
- [21] I. G. A. N. Lestari, N. M. R. M. Dewi, K. G. Meiliana, and I. K. A. A. Aryanto, “Effectiveness of AdaBoost and XGBoost Algorithms in Sentiment Analysis of Movie Reviews,” *J. Appl. Informatics Comput.*, vol. 9, no. 2, pp. 258–264, 2025, doi: 10.30871/jaic.v9i2.9077.
- [22] M. I. Anugrah, J. Zeniarja, and D. S. Setiawan, “Peningkatan Performa Model Hard Voting Classifier dengan Teknik Oversampling ADASYN pada Penyakit Diabetes,” *Edumatic J. Pendidik. Inform.*, vol. 8, no. 1, pp. 290–299, 2024, doi: 10.29408/edumatic.v8i1.25838.
- [23] C. Magnolia, A. Nurhopipah, and B. A. Kusuma, “Penanganan Imbalanced Dataset untuk Klasifikasi Komentar Program Kampus Merdeka Pada Aplikasi Twitter,” *Edu Komputika J.*, vol. 9, no. 2, pp. 105–113, 2022, doi: 10.15294/edukomputika.v9i2.61854.
- [24] N. F. Putri, M. F. Hidayattullah, and D. I. Af'idah, “Sentimen Analisis Kota Tegal Berbasis Aspek Menggunakan Algoritma Naïve Bayes,” *Infomatek*, vol. 26, no. 1, pp. 45–54, 2024, doi: 10.23969/infomatek.v26i1.11209.