

Optimasi Ulasan Palsu Menggunakan ADASYN Dan SMOTE

Aldrian Firmansyah Putranto¹, Agung Septian², Dicky Bagus Prasetyo³, RhoY Emiliano De Mozzagie⁴,
Rasyid Faiz Abdillah⁵

^{1,2,3,4,5}program Studi informatika, Fakultas Ilmu Komputer, Universitas Amikom Purwokerto, Purwokerto, Indonesia
surel: ¹ryury345@gmail.com, ²agung.septian.hartanto@gmail.com, ³batalion2725@gmail.com, ⁴rasyidfaizz1@gmail.com, ⁵rhoymozzagie@gmail.com

Info Artikel

Sejarah artikel:

Diterima 18-07-2025
Revisi tgl 19-08-2025
Diterima 28-08-2025

Kata kunci:

Fake review
ADASYN
SMOTE
machine learning
Yelp Dataset

ABSTRAK

Ulasan pengguna secara daring menjadi acuan penting dalam pengambilan keputusan konsumen di berbagai platform e-commerce. Maraknya ulasan palsu (fake review) yang sengaja dibuat untuk meningkatkan citra produk atau menjatuhkan pesaing menimbulkan tantangan serius dalam menjamin keaslian informasi. Penelitian ini bertujuan untuk meningkatkan akurasi deteksi ulasan palsu dengan menerapkan teknik oversampling lanjutan, yaitu SMOTE (Synthetic Minority Over-sampling Technique) dan ADASYN (Adaptive Synthetic Sampling) untuk menangani ketidakseimbangan data yang umum pada kasus ini. Dataset yang digunakan adalah Yelp Review Dataset dari Kaggle dengan jumlah data mencapai 50.000 ulasan yang telah melalui tahapan preprocessing seperti tokenisasi, normalisasi, stopword removal, dan stemming. Ekstraksi fitur dilakukan menggunakan metode TF-IDF. Model klasifikasi yang digunakan adalah XGBoost, LinearSVC, dan SGDClassifier. Hasil pengujian menunjukkan bahwa pendekatan SMOTE mampu meningkatkan performa deteksi kelas minoritas (ulasan palsu) dengan F1-score mencapai 0.75 meskipun mengorbankan akurasi kelas minoritas. Sementara itu, ADASYN belum menunjukkan performa optimal pada subset data kecil. Temuan ini mendukung hasil dari penelitian sebelumnya oleh Hameed et al. (2023) menunjukkan adanya potensi peningkatan performa melalui pendekatan balancing yang lebih adaptif. Penelitian ini memberikan kontribusi penting dalam pengembangan sistem deteksi ulasan palsu yang lebih seimbang dan andal.

Penulis yang sesuai:

Aldrian Firmansyah Putranto
Program Studi Informatika Fakultas Ilmu Komputer Universitas Amikom Purwokerto
Email: ryury345@gmail.com

1. PENDAHULUAN

Ulasan online telah menjadi bagian integral dalam proses pengambilan keputusan konsumen pada era digital. Konsumen kini semakin mengandalkan ulasan sebelum melakukan pembelian barang atau jasa sehingga ulasan memiliki pengaruh besar terhadap reputasi bisnis. Sayangnya, pengaruh ini juga dimanfaatkan oleh pihak tertentu dengan menyebarkan ulasan palsu (fake review) demi kepentingan tertentu baik untuk meningkatkan citra produk sendiri maupun untuk menjatuhkan pesaing. Hameed et al. (2023)

menyebu bahwa ulasan palsu dapat sangat merugikan pelaku usaha terutama bisnis kecil yang rentan terhadap manipulasi reputasi[1].

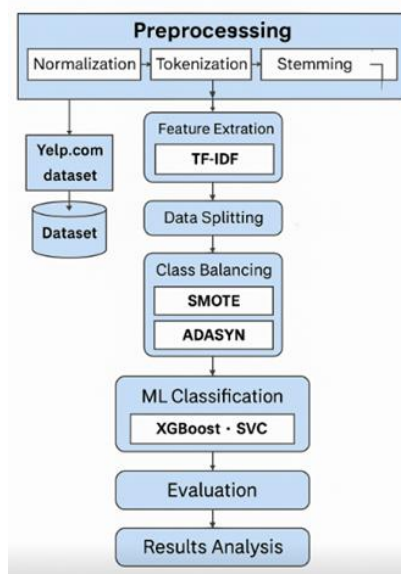
Deteksi ulasan palsu telah menjadi topik penting dalam bidang *natural language processing* (NLP) dan *machine learning*. Tantangan utama yang dihadapi dalam penelitian ini adalah ketidakseimbangan data (class imbalance) dimana ulasan asli (not spam) jauh lebih banyak dibandingkan ulasan palsu (spam). Ketidakseimbangan ini menyebabkan model klasifikasi cenderung bias terhadap kelas mayoritas, sehingga gagal mendeteksi kelas minoritas yang justru menjadi target utama dalam kasus ini[2].

Beberapa pendekatan balancing seperti random oversampling dan undersampling telah digunakan tetapi memiliki kelemahan tersendiri seperti overfitting atau hilangnya informasi penting. Untuk mengatasi hal ini, penelitian ini mengusulkan penggunaan dua teknik oversampling lanjutan, yaitu SMOTE dan ADASYN yang mampu menghasilkan data sintesis dengan cara yang lebih adaptif dan realistis. Penelitian ini juga menggunakan metode *term frequency-inverse document frequency* (TF-IDF) untuk ekstraksi fitur, serta menerapkan model XGBoost, LinearSVC, dan SGDClassifier untuk klasifikasi[3].

Penelitian ini bertujuan untuk mengevaluasi dan membandingkan efektivitas teknik SMOTE dan ADASYN dalam mendeteksi ulasan palsu pada data Yelp yang telah diseimbangkan serta melihat bagaimana performa model terhadap masing-masing pendekatan. Hasil dari penelitian ini diharapkan dapat memberikan kontribusi nyata dalam upaya pengembangan sistem deteksi ulasan palsu yang lebih akurat dan dapat diandalkan sekaligus memperluas temuan sebelumnya yang telah dilakukan oleh peneliti lain seperti Hameed et al. (2023)[4].

2. METODE

Bagian metode penelitian ini menjelaskan proses sistematis yang digunakan untuk melakukan penelitian. Gambar alir dari setiap tahapan proses dapat digunakan untuk menggambarkan secara visual dan terperinci alur kerja system. Gambar berikut menunjukkan flowchart yang menggambarkan proses utama, mulai dari pengumpulan data, pengolahan data, hingga hasil akhir dari sistem yang dikembangkan.



Gambar 1. Flowchart metode

2.1 Alur Metode

2.1.1. Dataset

Dataset yang digunakan berasal dari Kaggle. Dataset ini berisi ulasan pengguna, Data ini mencakup informasi tentang bisnis di 8 wilayah metropolitan di Amerika Serikat dan Kanada. Dataset ini terdiri dari lima berkas JSON. Banyak dataset tersebut ialah 6.990.280.

2.1.2. Preprocessing

Preprocessing merupakan tahap krusial dalam sistem deteksi ulasan palsu karena bertujuan untuk membersihkan dan menyiapkan teks agar dapat diproses secara efisien oleh algoritma machine learning. Teks ulasan pada dataset Yelp umumnya panjang, mengandung karakter khusus, variasi bentuk kata, dan kata-kata tidak penting (noise), sehingga diperlukan beberapa tahapan transformasi sebelum ekstraksi fitur. Langkah-langkah preprocessing yang diterapkan meliputi:

- a. Normalization
Proses ini bertujuan menyamakan format seluruh teks agar seragam. Langkah-langkah normalisasi yang dilakukan dengan Mengubah seluruh huruf menjadi huruf kecil (*lowercasing*): “Wow! This place is AMAZING!!! 10/10 would come again :) #happy” > “wow this place is amazing would come again happy”
- b. Tokenization
Teks dipecah menjadi token (kata) berdasarkan spasi. Tokenisasi memungkinkan setiap kata diproses sebagai fitur individual. Kalimat: " great food and service " Token hasil: ["great", "food", "and", "service"]
- c. Stemming
Mengubah kata ke bentuk dasar untuk menghindari redundansi fitur akibat variasi bentuk kata. Contoh: “running”, “ran”, “runs” > “run”

2.1.3. Feature Extraction (TF-IDF)

Setelah teks dibersihkan, kata-kata harus diubah ke dalam format numerik agar bisa diproses oleh model machine learning. TF-IDF (Term Frequency - Inverse Document Frequency) adalah metode untuk memberi bobot pada kata-kata berdasarkan: Seberapa sering kata muncul di dokumen tertentu (*TF*) dan Seberapa jarang kata itu muncul di seluruh kumpulan dokumen (*IDF*)

2.1.4. Data Splitting (Train & Test)

Data dibagi menjadi dua bagian yaitu training set 80% Digunakan untuk melatih model agar mampu mengenali pola dari data Real dan Fake. Testing set 20% untuk menguji performa model pada data yang belum pernah dilihat sebelumnya. Contoh: Jika dataset berisi 50.000 review, 40.000 review digunakan untuk pelatihan model dan 10.000 review digunakan untuk pengujian. Pembagian dilakukan menggunakan chunk split untuk memastikan proporsi kelas seimbang di kedua subset.

2.1.5. Class Balancing (SMOTE & ADASYN)

Dataset asli bersifat imbalanced, dengan jumlah review *Real* (label 0) lebih sedikit dibanding *Fake* (label 1). Ketidakseimbangan ini membuat model bias terhadap kelas mayoritas. Untuk mengatasinya gunakan dua teknik oversampling modern:

- a. SMOTE
SMOTE membuat data sintetis untuk kelas minoritas (*Real*) dengan cara interpolasi antar data tetangga terdekat. Hasilnya, Jumlah data *Real* = jumlah *Fake*. Distribusi kelas menjadi seimbang
- b. ADASYN
ADASYN juga membuat data sintetis, tetapi lebih adaptif, Menambahkan lebih banyak data sintetis di area yang sulit diprediksi (ambiguous) dan Fokus ke perbatasan decision boundary.

2.1.6. ML Classification (XGBoost, LinearSVC, SGDClassifier)

Tiga algoritma Machine Learning digunakan:

Tabel 1. Algoritma Machine Learning

Model	Deskripsi
XGBoost	Model ensemble berbasis pohon keputusan. Efisien dan akurat untuk data besar
LinearSVC	Support Vector Classifier linear. Cocok untuk data hasil TF-IDF
SGDClassifier	Model linear dengan pembaruan cepat menggunakan Stochastic Gradient Descent

Model dilatih menggunakan dataset hasil balancing SMOTE dan ADASYN, lalu diuji pada test set asli.

2.1.7. Perbandingan Data

Tabel 2. Perbandingan Data

Aspek	Dokumen	Jurnal
Nama Dataset	Yelp Dataset (Kaggle)	Yelp Dataset (Kaggle)
Ukuran Dataset Asli	± 6.990.280 review (5GB total)	Tidak disebut
Jumlah Data yang Digunakan	Hanya 50.000 data pertama, dipilih secara acak	±67.722 review (60.929 not spam + 6.793 spam)
Alasan Pembatasan Data	Untuk menghindari lag memori dan mempercepat proses eksperimen	Tidak dijelaskan
Metode Labeling	Label 1 = rating 5 (Fake), 0 = rating 3-4 (Real)	Label 1 = Spam, 0 = Not Spam (tapi cara penentuan tidak dijelaskan)
Preprocessing	Lowercasing, tokenization, stopwords removal, stemming, TF-IDF	Tokenization, stemming, stopwords removal, TF-IDF
Teknik Balancing	SMOTE & ADASYN	Hanya random oversampling & undersampling

Dataset yang digunakan dalam jurnal PDF terdiri dari sekitar 67.722 review, dengan dominasi kelas Not Spam sebanyak 60.929 dan hanya 6.793 Spam. Sebaliknya, pada penelitian ini hanya digunakan 50.000 data yang diambil dari dataset Yelp dengan proporsi yang lebih seimbang antara kelas Fake dan Real. Pembatasan ini dilakukan karena keterbatasan memori dan efisiensi pelatihan. Selain itu, metode labeling dan atribut yang digunakan dalam proyek ini lebih transparan dan bervariasi dibandingkan jurnal, yang tidak menjelaskan secara rinci asal-usul label maupun struktur fitur.

2.1.8. Atribut Data

Tabel 3. Atribut Data

Nama Kolom	Tipe Data	Penjelasan
review_id	string	ID unik dari ulasan. Setiap ulasan memiliki ID yang berbeda
user_id	string	ID pengguna yang menulis ulasan. Dapat digunakan untuk menggabungkan data dari user.json.
business_id	string	ID bisnis yang diulas. Digunakan untuk menggabungkan dengan data bisnis terkait
stars	integer	Rating yang diberikan oleh pengguna terhadap bisnis (1-5)
text	string	Isi lengkap dari ulasan. Digunakan dalam analisis sentimen dan NLP
useful	string	Jumlah pengguna yang menandai ulasan sebagai berguna
funny	string	Jumlah pengguna yang menandai ulasan sebagai lucu
cool	string	Jumlah pengguna yang menandai ulasan sebagai menarik atau keren
sentiment	string	Nilai sentimen hasil analisis terhadap teks review (positif atau negatif)
label	integer	Target klasifikasi, 0 untuk Not Spam dan 1 untuk Spam

2. HASIL DAN PEMBAHASAN

Setelah model membuat prediksi, performanya diukur menggunakan beberapa metrik accuracy yaitu Persentase prediksi yang benar dari seluruh data uji, precision seberapa tepat model saat memprediksi review sebagai fake, recall: Seberapa banyak review palsu yang berhasil dideteksi, F1-Score: Harmoni antara precision dan recall. Dan hasilnya terdapat pada tabel 4 dan 5

Tabel 4. Hasil SMOTE

Model	Real (F1)	Fake (F1)	Analisis
XGBoost	0.00	0.75	Hanya mengenali Fake Review
LinearSVC	0.00	0.75	Sama seperti XGBoost
SGDClassifier	0.57	0.00	Hanya mengenali Real Review

Tabel 5. Hasil ADASYN

Model	Real (F1)	Fake (F1)	Analisis
XGBoost	0.57	0.00	Hanya mengenali Real Review
LinearSVC	0.57	0.00	Sama seperti XGBoost
SGDClassifier	0.57	0.00	Konsisten hanya memprediksi kelas Real



4. Keterbatasan

4.1. Unbalanced Data – Model Sangat Bias

- Dataset asli sangat tidak seimbang: mayoritas review adalah not spam (kelas 0), sedangkan spam (kelas 1) sangat sedikit.
- Model seperti XGBoost, SVC, dan SGD hanya belajar mengenali kelas mayoritas.
- F1-Score dan Recall untuk kelas spam = 0.00 — ini berarti model gagal sepenuhnya mengenali review palsu, meski akurasi keseluruhan terlihat tinggi (90%).

4.2. Balanced Data – Performa Merata Tapi Turun

- Dataset diseimbangkan, dan model diuji ulang.
- Hasil lebih seimbang untuk kedua kelas (spam dan not spam), tapi akurasi dan F1-score turun drastis (sekitar 60%).

4.3. Under-sampling – Menghilangkan Informasi Penting

- Jumlah data mayoritas dikurangi agar setara dengan data minoritas.
- Recall kelas spam meningkat (hingga 0.72), tapi akurasi dan precision menurun parah.

4.4. Random Oversampling – Rentan Overfitting

- Data minoritas (spam) digandakan secara acak agar jumlahnya sama dengan kelas mayoritas.
- Recall kelas spam sedikit meningkat, tapi precision tetap rendah (~0.14).

4.5. Tidak Menggunakan Teknik Sampling Lanjutan

Jurnal tidak mencoba metode sampling yang lebih canggih, seperti:

- SMOTE (Synthetic Minority Over-sampling Technique)
- ADASYN

Sumber Dataset: [dataset](#)

5. KESIMPULAN

Berdasarkan eksperimen yang dilakukan, diperoleh beberapa kesimpulan utama: SMOTE menghasilkan F1-score terbaik (0.75) pada kelas Fake, khususnya pada model XGBoost dan LinearSVC. Ini menunjukkan kemampuannya dalam membantu model mengenali review palsu. Namun, hal ini diikuti dengan kegagalan total mendeteksi review Real ($F1 = 0.00$), menunjukkan adanya overfitting ke kelas minoritas. Model SGDClassifier menunjukkan performa yang kontras: hanya mengenali kelas mayoritas (baik di SMOTE maupun ADASYN), sehingga tidak seimbang dan gagal menangkap kedua kelas. Dibandingkan dengan teknik random oversampling dan undersampling seperti pada jurnal, SMOTE terbukti memberikan peningkatan yang signifikan pada kelas minoritas (Fake), meskipun dengan trade-off terhadap kelas mayoritas. Penggunaan ADASYN ini belum menghasilkan performa optimal ($F1\text{-score} = 0.00$ untuk kelas Fake). Hal ini kemungkinan besar disebabkan oleh pembatasan jumlah data yang digunakan dalam pelatihan, yaitu hanya sebanyak 50.000 data dalam dataset Yelp. Karena ADASYN bekerja dengan prinsip adaptif berdasarkan densitas lokal dari kelas minoritas, jumlah data yang terlalu kecil dapat menyebabkan fitur minoritas tidak terwakili secara memadai di ruang vektorisasi. Oleh karena itu, dalam skenario data penuh, ADASYN berpotensi menunjukkan performa yang lebih baik.

Saran, perlu ditegaskan bahwa penerapan teknik SMOTE dianggap membantu skor F1 kelas minoritas (Palsu), dan pada saat yang sama, menyebabkan overfitting kelas tersebut sementara mengabaikan kelas mayoritas (Nyata). Modifikasi model akan menjadi vital untuk memungkinkan penggabungan regularisasi atau penentuan bobot kelas yang lebih seimbang antara kelas minoritas dan mayoritas. Pengaturan semacam itu akan mendukung pengenalan pola pada kedua kelas dan mencegahnya menjadi timpang dan menguntungkan beberapa kelas. Selain itu, hasil yang baik dapat diperoleh dengan menggunakan SMOTE, tetapi metode resampling lain, seperti NearMiss atau ClusterCentroids, perlu dievaluasi; hal ini dapat memberikan

keseimbangan yang lebih baik antara kelas-kelas yang dimaksud, yang pada gilirannya akan memungkinkan model untuk mempelajari sesuatu dari kedua kelas.

Dalam kasus ADASYN, kinerja kelas palsu yang rendah disebabkan oleh kurangnya dataset. Oleh karena itu, langkah ini akan meningkatkan dataset pelatihan. Hal ini jelas akan memungkinkan model untuk menangani distribusi kelas minoritas dan mayoritas dengan lebih baik. Selain itu, penggunaan teknik ensemble seperti tumpukan atau pengklasifikasi voting menggunakan berbagai model juga dapat meningkatkan stabilitas dan akurasi model karena dapat memanfaatkan kekuatan yang berbeda dari berbagai algoritma saat memprediksi kelas. Algoritma lain seperti Random Forest atau LightGBM juga dapat dicoba untuk kinerja yang lebih baik karena model saat ini mungkin sangat bias terhadap kelas mayoritas.

Metode ekspansi data di atas dapat mencakup injeksi derau atau penghapusan acak data kelas minoritas untuk meningkatkan variasi dalam data pelatihan guna mendorong generalisasi model, sehingga mencegahnya terkunci pada pola tertentu. Selain skor F1, penggabungan matriks kebingungan akan memberikan gambaran model yang lebih lengkap. Hal ini juga akan menunjukkan kesalahan spesifik apa yang terjadi di setiap kelas, sehingga memberikan pemahaman yang lebih mendasar untuk perbaikan lebih lanjut.

REFERENSI

- [1] M. Sokolova and G. Lapalme, "A systematic analysis of performance measures for classification tasks," *Inf. Process. Manag.*, vol. 45, no. 4, pp. 427–437, Jul. 2009, doi: [10.1016/j.ipm.2009.03.002](https://doi.org/10.1016/j.ipm.2009.03.002).
- [2] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognit. Lett.*, vol. 27, no. 8, pp. 861–874, Jun. 2006, doi: [10.1016/j.patrec.2005.10.010](https://doi.org/10.1016/j.patrec.2005.10.010).
- [3] Y. Maraden *et al.*, "Enhancing electricity theft detection through K-nearest neighbors and logistic regression algorithms with synthetic minority oversampling technique: A case study on State Electricity Company (PLN) customer data," *Energies*, vol. 16, no. 14, p. 5405, Jul. 2023, doi: [10.3390/en16145405](https://doi.org/10.3390/en16145405).
- [4] N. Ali, D. Neagu, and P. Trundle, "Evaluation of k-nearest neighbour classifier performance for heterogeneous data sets," *SN Appl. Sci.*, vol. 1, no. 12, Nov. 2019, doi: [10.1007/s42452-019-1356-9](https://doi.org/10.1007/s42452-019-1356-9).
- [5] S. Dreiseitl and L. Ohno-Machado, "Logistic regression and artificial neural network classification models: A methodology review," *J. Biomed. Inform.*, vol. 35, no. 5–6, pp. 352–359, Oct. 2002, doi: [10.1016/s1532-0464\(03\)00034-0](https://doi.org/10.1016/s1532-0464(03)00034-0).
- [6] Y. Rimal *et al.*, "Comparative analysis of heart disease prediction using logistic regression, SVM, KNN, and random forest with cross-validation for improved accuracy," *Sci. Rep.*, vol. 15, no. 1, Apr. 2025, doi: [10.1038/s41598-025-93675-1](https://doi.org/10.1038/s41598-025-93675-1).
- [7] F. Y. A'la, "Optimasi klasifikasi sentimen ulasan game berbahasa Indonesia: IndoBERT dan SMOTE untuk menangani ketidakseimbangan kelas," *Edumatic: J. Pendidik. Inform.*, vol. 9, no. 1, pp. 256–265, Apr. 2025, doi: [10.29408/edumatic.v9i1.29666](https://doi.org/10.29408/edumatic.v9i1.29666).
- [8] A. Firmansyah and A. Aziz, "Optimasi K-nearest neighbor menggunakan algoritma SMOTE untuk mengatasi imbalance class pada klasifikasi analisis sentimen," *J. Mahasiswa Tek. Inform.*, vol. 7, no. 6, 2023.
- [9] W. Hidayat, M. Ardiansyah, and A. Setyanto, "Pengaruh algoritma ADASYN dan SMOTE terhadap performa Support Vector Machine pada ketidakseimbangan dataset Airbnb," *Edumatic: J. Pendidik. Inform.*, vol. 5, no. 1, pp. 11–20, Jun. 2021, doi: [10.29408/edumatic.v5i1.3125](https://doi.org/10.29408/edumatic.v5i1.3125).
- [10] N. K. Goyal, K. Gupta, D. Goyal, and A. Pal, "Fake review detection using ensemble techniques by the fusion of chronology, aspect and sentiment of reviews and oversampling by SMOTE," 2023.
- [11] J. Beinecke and D. Heider, "Gaussian noise up-sampling is better suited than SMOTE and ADASYN for clinical decision making," *BioData Min.*, vol. 14, no. 1, Nov. 2021, doi: [10.1186/s13040-021-00283-6](https://doi.org/10.1186/s13040-021-00283-6).
- [12] F. Last, G. Douzas, and F. Bacao, "Oversampling for imbalanced learning based on K-means and SMOTE."
- [13] A. M. Halim, M. Dwifabri, and F. Nhita, "Handling imbalanced data sets using SMOTE and ADASYN to improve classification performance of Ecoli data sets," *Build. Informatics, Technol. Sci. (BITS)*, vol. 5, no. 1, Jun. 2023, doi: [10.47065/bits.v5i1.3647](https://doi.org/10.47065/bits.v5i1.3647).
- [14] C. Kim and H. G. Kim, "Efficient detection of irrelevant user reviews using machine learning," *Appl. Sci.*, vol. 14, no. 16, p. 6900, Aug. 2024, doi: [10.3390/app14166900](https://doi.org/10.3390/app14166900).
- [15] J. Hassannataj-Joloudari, A. Marefat, M. A. Nematollahi, S. S. Oyelere, and S. Hussain, "Effective class-imbalance learning based on SMOTE and convolutional neural networks," *Appl. Sci.*, vol. 13, no. 6, p. 4006, Mar. 2023, doi: [10.3390/app13064006](https://doi.org/10.3390/app13064006).

