

Optimalisasi Model Klasifikasi Sentimen Netizen Terhadap Merek Tas Luar Negeri

Mochamad Nurul Huda^{1✉}, Daffa Almer Fauzan², Muhammad Raihan Satrio Putra Pamungkas³, Nikita Sabila Ratnadewi⁴, Azzahra Ayu Vahendra⁵

^{1,2,3,4,5} Program Studi Rekayasa Perangkat Lunak, Universitas Pendidikan Indonesia

mochamadnuru1huda16@upi.edu

Abstract

Research on text mining has grown more than ever in various sectors. Public figures have also grown in interest towards the field and have the tendency to get to know more about consumers' perceptions toward relevant goods and the reputation of an individual in social media. Sentiment analysis is a state-of-the-art technique that can be utilized to evaluate such trends or general views, for instance the reputation of a fashion brand. The dataset is built upon the crawled tweets that are relevant with the required topics which have the purpose to analyze the preferred fashion brand of the public. This study shows that the public leads to a positive notion toward foreign bag brands. The algorithms that are being compared includes Logistic Regression, Multinomial Naïve Bayes, Decision Tree, K-Nearest Neighbors, Random Forest, and Support Vector Machine. Support Vector Machine provides the best model which reaches 69% in accuracy. The Synthetic Minority Oversampling Technique (SMOTE) was also conducted to improve the model. Result shows that the Support Vector Machine model has successfully increased its accuracy by 13%, reaching an accuracy of 82%. The findings of this study can serve as a reference for future similar research on sentiment analysis, particularly in the development of datasets and classification models. Additionally, the SMOTE technique in this study has been proven to effectively mitigate class imbalance issues in the dataset.

Keywords: Sentiment Analysis, Brand, Machine Learning, Classification, SMOTE

Abstrak

Penelitian mengenai *text mining* telah mengalami peningkatan dibanding sebelumnya di dalam berbagai sektor. Figur publik juga semakin tertarik terhadap bidang tersebut dan memiliki kecenderungan untuk mengetahui lebih banyak mengenai persepsi konsumen terhadap suatu barang dan mengenai reputasi seseorang di media sosial. Analisis Sentimen merupakan sebuah teknik *state-of-the-art* yang dapat digunakan untuk mengevaluasi suatu tren atau pandangan umum mengenai suatu hal, misalnya reputasi sebuah merek *fashion*. Sumber himpunan data yang digunakan pada penelitian ini dibuat berdasarkan *crawling tweet* yang relevan dengan topik yang dibutuhkan, yang bertujuan untuk menganalisis merek *fashion* yang disukai oleh masyarakat. Penelitian ini menunjukkan bahwa persepsi masyarakat mengarah pada persepsi positif terhadap merek tas luar negeri. Pada penelitian ini, beberapa algoritma digunakan sebagai perbandingan, antara lain *Logistic Regression*, *Multinomial Naïve Bayes*, *Decision Tree*, *K-Nearest Neighbors*, *Random Forest*, dan *Support Vector Machine*. Hasil pengujian model menunjukkan algoritma *Support Vector Machine* memiliki performa terbaik dengan *accuracy* sebesar 69%. Kemudian digunakan teknik *Synthetic Minority Oversampling Technique* (SMOTE) untuk meningkatkan performa dari model. Hasil menunjukkan bahwa model algoritma *Support Vector Machine* telah berhasil ditingkatkan dengan *accuracy* sebesar 13%, mencapai *accuracy* sebesar 82%. Adapun hasil dari penelitian ini dapat dijadikan sebagai acuan untuk penelitian sejenis di masa depan mengenai analisis sentimen, khususnya dalam pembuatan *dataset* dan pengembangan model klasifikasi. Selain itu, penggunaan teknik SMOTE dalam penelitian ini telah terbukti dapat digunakan untuk mengatasi masalah ketidakseimbangan kelas pada *dataset*.

Kata kunci: Analisis Sentimen, Merek, Pembelajaran Mesin, Klasifikasi, SMOTE

KomtekInfo is licensed under a Creative Commons Attribution-Share Alike 4.0 International License.



1. Pendahuluan

Berbagai jenis *fashion* yang diciptakan di luar negeri ternyata memiliki atensi yang cukup tinggi bagi masyarakat Indonesia [1], salah satu contohnya adalah jenis tas. Saat ini telah ada opini yang menetap pada pola pikir masyarakat sehingga beranggapan bahwa semakin tinggi nilai tas yang dimiliki maka semakin tinggi pandangan orang lain terhadap pemilik tas tersebut [2]. Kedua hal tersebut menjadi indikator daya minat masyarakat terhadap *fashion* luar negeri cukup tinggi.

Telah menjadi hal yang umum bagi masyarakat atau konsumen sebelum melakukan pembelian barang, maka terlebih dahulu mereka akan mengidentifikasi dan membaca *review* serta umpan balik dari konsumen lainnya [3]. Terlebih sebuah studi menunjukkan dan memperkuat adanya pengaruh *review* dari konsumen ataupun seorang *influencer* terhadap minat untuk membeli barang [4,5]. Selain itu juga, *review* dan umpan balik konsumen memiliki urgensi tinggi bagi pemilik bisnis atau perusahaan, melalui proses identifikasi serta analisis bisnis yang dilakukan maka kelemahan produk dari sudut pandang konsumen dapat diketahui [6].

Pada sebuah studi menunjukkan bahwa review pada produk yang bersifat negatif memiliki dampak negatif terhadap keputusan konsumen untuk membeli suatu produk. Konsumen cenderung menghindari keputusan berisiko dan konsumen lebih sensitif terhadap kerugian daripada keuntungan. Review yang negatif menyebabkan peningkatan persepsi tentang kualitas produk yang buruk dan penurunan persepsi tentang nilai produk. Sehingga memiliki efek negatif yang signifikan terhadap persepsi harga dan keputusan pembelian [7].

Saat ini proses identifikasi review konsumen menggunakan pendekatan analisis sentimen dan algoritma pembelajaran mesin telah banyak digunakan dalam beberapa penelitian sebelumnya, seperti pada studi kasus perusahaan *start-up* di Indonesia [8,9]. Analisis sentimen memiliki kemudahan terutama dalam menghimpun dan menganalisis data secara *realtime*, sehingga hal tersebut menarik minat para peneliti maupun profesional di kalangan industri [10].

Tidak banyak penelitian terkait mengenai analisis sentimen pada lingkup pakaian atau *fashion* yang telah ditemukan. Penelitian pertama menggunakan *dataset* dari *Amazon Review* menggunakan algoritma *Support Vector Machine* memperoleh akurasi terbaiknya sebesar 88.0% dengan jumlah *dataset* seimbang. Namun label sentimen hanya terbatas pada 50,000 review positif dan 50,000 negatif [11].

Penelitian kedua yang dilakukan menggunakan *dataset* berasal dari proses *scrapping* melalui Twitter API sebanyak 9,652 *tweet*. Proses pemilihan algoritma dan model klasifikasi terbaik dilakukan secara otomatis menggunakan alat *Wolfram Matematika* [12]. Sehingga, pada penelitian tersebut akurasi kinerja model tidak diketahui secara eksplisit.

Penelitian ketiga menggunakan *dataset* pakaian dari situs *e-commerce* sebanyak 23,468 data dengan pemberian label positif serta negatif. Model selanjutnya dikembangkan menggunakan *hyperparameter* ukuran *batch* sebanyak 16, *learning rate* sebesar 0,001 dan *epoch* sebanyak 5 sehingga diperoleh akurasi validasi sebesar 81.5% [13].

Berdasarkan literatur belum adanya penelitian yang membahas secara spesifik pada merek tas. Serta belum adanya penelitian yang mengangkat optimalisasi ataupun improvisasi model prediksi menggunakan

suatu teknik seperti salah satu contohnya adalah SMOTE (*Synthetic Minority Over-sampling Technique*). Padahal sebuah studi menunjukkan teknik SMOTE meningkatkan kinerja model terutama untuk kasus *dataset* yang tidak seimbang [14].

Oleh karena itu, penelitian ini memiliki tujuan untuk membuat *dataset* baru yang berisi sentimen netizen mengenai beberapa merek tas, menganalisis *dataset* yang telah dibuat, dan mengembangkan model prediksi klasifikasi sentimen netizen terhadap merek tas luar negeri serta melakukan proses optimalisasi terhadap model tersebut menggunakan teknik SMOTE. Label yang dicantumkan pada model klasifikasi antara lain label positif, negatif dan netral. Pembuatan model klasifikasi bertujuan untuk memprediksi sentimen netizen di masa depan, apakah sentimen tersebut termasuk ke dalam kategori positif, negatif, atau netral.

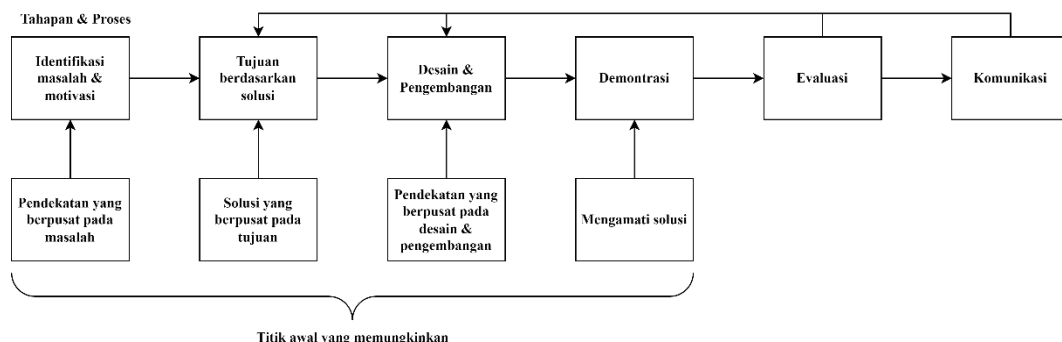
Luaran yang dihasilkan yaitu model klasifikasi menggunakan beragam jenis algoritma serta model yang telah ter optimalisasi menggunakan teknik SMOTE. Hasil dari penelitian ini dapat dijadikan sebagai acuan untuk penelitian selanjutnya mengenai analisis sentimen, khususnya dalam pembuatan *dataset* dan pengembangan model klasifikasi. Selain itu, penggunaan teknik SMOTE pada penelitian ini, dapat digunakan sebagai acuan untuk mengatasi masalah ketidakseimbangan kelas atau label pada *dataset*.

2. Metodologi Penelitian.

2.1. Desain Penelitian

Penelitian ini dilakukan dengan menggunakan *Design Science Research Process* (DSRP) yang ditunjukkan oleh Gambar 1 yang terdiri atas identifikasi masalah dan motivasi, tujuan dari sebuah solusi, desain dan pengembangan, demonstrasi, evaluasi dan komunikasi [15]. Desain penelitian ini dipilih karena telah banyak digunakan dalam penelitian rumpun ilmu komputer dan sesuai dengan penelitian yang dilakukan.

Pada Gambar 1 diketahui urutan langkah-langkah yang dilakukan dalam penelitian ini. Tahap pertama yaitu mengidentifikasi masalah yang ingin dipecahkan atau kebutuhan yang harus dipenuhi. Kemudian menetapkan tujuan dari solusi yang ingin dihasilkan. Selanjutnya merancang dan mengembangkan solusi yang akan membantu dalam menyelesaikan masalah. Solusi



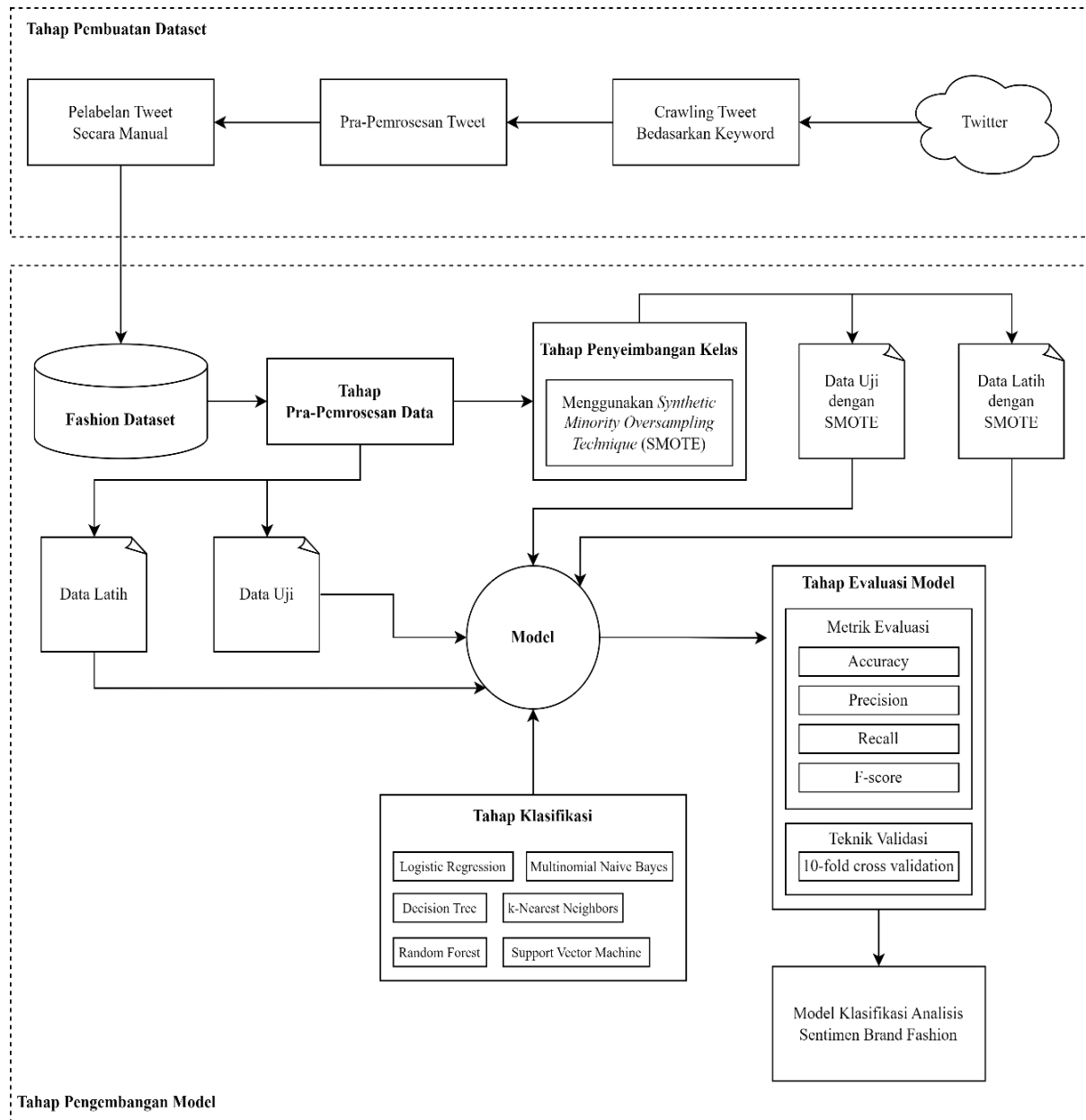
Gambar 1. Desain Penelitian

tersebut dapat berupa teknologi baru dan kerangka kerja yang berbeda.

Tahap selanjutnya yaitu mendemonstrasikan solusi yang telah dibuat ke dalam bentuk prototipe atau model. Kemudian mengevaluasi keefektifan dari solusi yang telah dibuat. Terakhir, mengkomunikasikan hasil penelitian dengan cara mempublikasikan hasil penelitian ke dalam bentuk artikel ilmiah.

2.2. Kerangka Kerja Penelitian

Setelah mengidentifikasi masalah dan tujuan, tahap selanjutnya adalah desain dan pengembangan. Untuk melakukan tahap tersebut hingga tahap akhir, diusulkan kerangka kerja yang terdiri dari pembuatan *dataset* dan pengembangan model yang terdapat pada Gambar 2. Secara garis besar kerangka kerja tersebut terdiri atas dua tahap, yaitu tahap pembuatan *dataset* dan tahap pengembangan model pembelajaran mesin.



Gambar 2. Kerangka Kerja yang Diusulkan

Pada tahap pembuatan *dataset*, langkah pertama yang dilakukan adalah pengambilan data dengan cara *crawling tweet* pada Twitter berdasarkan kata kunci sesuai dengan kebutuhan. Setelah data tersebut didapatkan, kemudian dilakukan pra-pemrosesan data untuk membersihkan data sebelum digunakan, dan

terakhir melakukan pelabelan sentimen pada data *tweet* secara manual dengan memperhatikan konteks *tweet* tersebut.

Pada tahap pengembangan model, langkah pertama yang dilakukan adalah pra-pemrosesan data pada *dataset* yang telah dibuat sebelumnya. Pra-pemrosesan

data di sini merupakan tahapan persiapan data yang bertujuan untuk membersihkan data dan menyesuaikan data sebelum dianalisis oleh model pembelajaran mesin. Setelah pra-pemrosesan data selesai, langkah selanjutnya terbagi menjadi dua eksperimen. Eksperimen yang pertama langsung kepada tahapan membagi *dataset* menjadi data latih dan data uji. Sedangkan eksperimen kedua, melakukan teknik SMOTE setelah pra-pemrosesan data, dan kemudian membagi *dataset* menjadi data latih dan data uji.

Selanjutnya data-data ini digunakan dalam pembuatan model dengan menggunakan beberapa algoritma untuk permasalahan klasifikasi seperti *Logistic Regression*, *Multinomial Naïve Bayes*, *Decision Tree*, *K-Nearest Neighbors*, *Random Forest*, dan *Support Vector Machine*. Kemudian untuk mengevaluasi model pembelajaran mesin, digunakan beberapa metrik evaluasi yang umum untuk masalah klasifikasi, seperti *Accuracy*, *Precision*, *Recall*, dan *F-score*. Selain itu, pada saat proses pembuatan model digunakan teknik *10-fold cross validation* untuk mengevaluasi model.

2.3. Sumber Himpunan Data

Sumber himpunan data yang digunakan dalam penelitian ini adalah data *tweet* dari Twitter berdasarkan beberapa kata kunci yang telah ditetapkan sebelumnya. Data *tweet* ini di ambil dengan menggunakan Twitter API yang digunakan untuk mendapatkan data dari platform media sosial Twitter. Adapun beberapa kata kunci yang digunakan, dan jumlah data *tweet* yang didapatkan pada saat pengambilan data, terdapat pada Tabel 1.

Tabel 1. Hasil Crawling Twitter

No.	Nama Merek	Kata Kunci	Jumlah Data
1	Gucci	gucci bag	821
2	Chanel	chanel bag	698
3	Dior	dior bag	435
4	Louis Vuitton	louis vuitton bag	359
5	Prada	prada bag	336
6	Hermes	hermes bag	241
7	Supreme	supreme bag	220
			3110

Berdasarkan hasil *crawling* Twitter pada Tabel 1, didapatkan data dengan jumlah 3110 *tweet* dari beberapa kata kunci merek tas terkenal dunia. Tabel 1 menunjukkan pengurutan setiap data dilakukan berdasarkan jumlah data terbanyak. Data terbanyak diraih oleh kata kunci “*gucci bag*” dan terendah diraih oleh kata kunci “*supreme bag*”. Setelah data terhimpun, kemudian dilakukan proses pra-pemrosesan data untuk menghilangkan data duplikat sehingga keseluruhan data yang digunakan pada penelitian ini adalah berjumlah 2881 *tweet* [16].

2.4. Teknik yang Digunakan

Teknik pertama yang digunakan adalah pra-pemrosesan data, spesifiknya terhadap data tekstual. Pra-pemrosesan tersebut dilakukan agar data yang

digunakan memiliki format yang bersih, konsisten, dan siap untuk dianalisis lebih lanjut dengan harapan memberikan performa yang lebih baik secara signifikan [17]. Langkah-langkah yang dilakukan secara berturut-turut mencakup *case folding*, *link removal*, *punctuation removal*, *tokenization*, *stopwords removal*, dan *stemming*.

Teknik kedua yang digunakan adalah *Count vectorizer* yang merupakan teknik *feature extraction* dan berperan dalam menggambarkan koleksi kata dalam bentuk matriks-matriks [18]. Penelitian berkaitan dengan klasifikasi teks condong menggunakan teknik ini, sebagaimana telah digunakan dalam penelitian sebelumnya mengenai klasifikasi sentimen pesan berindikasi *cyberbullying* [19].

Teknik terakhir yang diimplementasikan dalam penelitian ini berperan penting sebagai teknik optimalisasi yaitu teknik *oversampling* bernama SMOTE, teknik ini *mempopulasikan* kelas *dataset* minoritas dengan sampel data sintetis dari kelas tersebut menggunakan pendekatan *K-Nearest Neighbors* [20]. Teknik ini digunakan untuk mengatasi data yang tidak seimbang. Meskipun sebuah studi menunjukkan bahwa secara inferensial kinerja model dari *dataset* seimbang dan tidak seimbang tidak memiliki perbedaan signifikan [21]. Tetapi pada penelitian ini berupaya untuk menelusuri kinerja terbaiknya dari setiap model.

2.5. Parameter Evaluasi Kinerja

Penelitian ini menggunakan teknik klasifikasi dalam pembuatan model, sehingga digunakan metrik-metrik evaluasi yang umum digunakan untuk masalah klasifikasi, yaitu *Accuracy*, *Precision*, *Recall*, dan *F-Score* [22]. *Accuracy* mengukur proporsi antara hasil klasifikasi data yang benar dengan seluruh populasi. Rumus untuk menentukan *accuracy* ini disajikan pada Persamaan 1.

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

Precision mengukur rasio ketepatan hasil dalam *output* sistem. Metrik ini mengindikasikan bahwa semakin banyak jumlah hasil yang benar dan semakin sedikit jumlah hasil yang salah, maka semakin tinggi *precision* dari operasi yang dilakukan. Rumus untuk menentukan *precision* ini disajikan pada Persamaan 2.

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

Recall mengukur rasio hasil yang benar yang diberikan sistem dibandingkan dengan seluruh hasil. Metrik ini mengindikasikan bahwa semakin banyak jumlah hasil yang benar dan semakin sedikit jumlah hasil yang benar yang didiagnosis secara salah, maka semakin tinggi kemampuan *recall* pada operasi yang dilakukan. Rumus untuk menentukan *recall* ini disajikan pada Persamaan 3.

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

Rumus *F-score* merupakan kombinasi dari *precision* dan *recall*. Metrik ini mengindikasikan jumlah *precision* dan kemampuan *recall* dari operasi yang dilakukan. Rumus untuk menentukan *F-score* ini disajikan pada Persamaan 4.

$$F - score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (4)$$

Semua rumus-rumus tersebut mengacu pada *confusion matrix* pada Tabel 2. *Confusion matrix* adalah tabel yang digunakan untuk menghitung jumlah data yang diklasifikasikan dengan benar atau salah. *Confusion matrix* memiliki empat elemen yaitu *True Positive* (TP), *False Positive* (FP), *True Negative* (TN), dan *False Negative* (FN).

Tabel 2. *Confusion Matrix*

Prediksi	Aktual	
	TP (<i>True Positive</i>)	FP (<i>False Positive</i>)
	TN (<i>True Negative</i>)	FN (<i>False Negative</i>)

3. Hasil dan Pembahasan

Pada bagian ini, dijelaskan tahap-tahap yang telah dilakukan, yang sebelumnya telah direncanakan dalam kerangka kerja penelitian. Secara garis besar, tahap yang dilakukan adalah pembuatan sumber himpunan data atau *dataset* dan pengembangan model klasifikasi. Dalam pembuatan sumber himpunan data menjelaskan mengenai proses pembuatan, jumlah data, dan hasil analisis sentimennya. Kemudian dalam pengembangan model klasifikasi menjelaskan proses pembuatan model, performa yang didapatkan dari setiap algoritma dan perbaikan performa yang didapatkan dari setiap model.

3.1. Pembuatan Sumber Himpunan Data

Setelah melakukan pelabelan sentimen secara manual pada sumber himpunan data atau *dataset* yang berjumlah 2881 *tweet* didapatkan hasil keseluruhan yang di perlihatkan pada Tabel 3. Sentimen yang diberikan di setiap *tweet* dikategorikan menjadi tiga, yaitu: netral, positif dan negatif. Pada Tabel 3 diketahui bahwa untuk sentimen positif berjumlah 1083 data, negatif 374 data, dan netral 1424 data. Setelah data sentimen diolah dengan beberapa teknik pra-pemrosesan data, kemudian di dapatkan bentuk *word cloud* untuk setiap jenis sentimennya, yang ditunjukkan oleh Gambar 3, Gambar 4, dan Gambar 5.

Tabel 3. Jenis Dataset

Sentimen	Jumlah
Positif	1083
Negatif	374
Netral	1424

Word cloud adalah grup atau kumpulan kata yang ditampilkan dalam berbagai ukuran [23]. Kata yang lebih diucapkan akan berukuran besar dari yang lain jika kata tersebut semakin banyak muncul di suatu data. *Word cloud* efektif untuk mengilustrasikan pendapat tentang subjek tertentu. Maka dari itu, *word cloud* digunakan pada pengujian ini. Dari hasil *word cloud* yang telah diperoleh dan diklasifikasi berdasarkan sentimen, telah dihasilkan beberapa kata yang sering muncul. Kata-kata tersebut merupakan representasi dari setiap sentimen sehingga dapat dianalisis untuk mendapatkan hasil kesimpulan.

Pada Gambar 3 diketahui bahwa sentimen netral merupakan kategori *tweet* yang tidak relevan terhadap kata kunci yang di ambil atau topik yang di bahas. Selain itu *tweet* yang berupa promosi juga termasuk ke dalam sentimen netral, oleh karena itu dalam *word cloud* banyak ditemukan kata *ebay* atau sebuah *marketplace*.

Gambar 3. *Word Cloud* Netral

Pada Gambar 4 diketahui bahwa sentimen positif merupakan kategori *tweet* yang relevan dan bersifat positif. Bersifat positif dalam artian bahwa *tweet* tersebut memiliki konteks yang menyatakan suka, memuji-muji dan ingin mempunyai produk dari suatu merek. Hasil dari *word cloud* menunjukkan bahwa merek Gucci berada pada urutan pertama, kemudian Chanel, Dior, Prada, Louis Vuitton, dan Hermes.

Gambar 4. *Word Cloud* Positif

Pada Gambar 5 diketahui bahwa sentimen negatif merupakan kategori *tweet* yang relevan dan bersifat negatif. Bersifat negatif di sini berarti bahwa *tweet* memiliki konteks yang menyatakan rasa tidak suka, benci, dan sarkasme terhadap suatu merek. Hasil dari

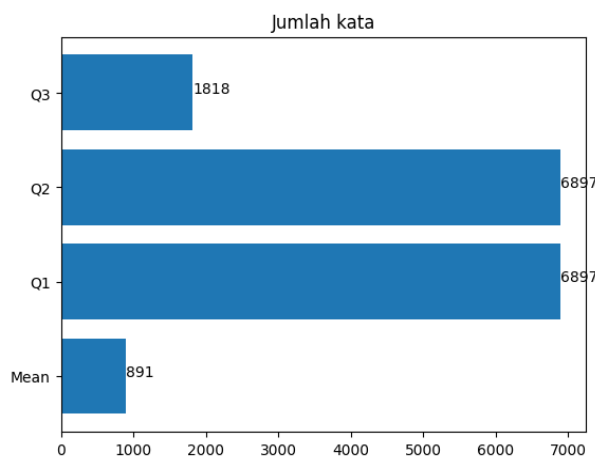
word cloud menunjukkan bahwa merek Gucci pada urutan pertama, kemudian Chanel, Prada, Louis Vuitton, Dior, dan Hermes.



Gambar 5. Word Cloud Negatif

3.2. Pengembangan Model Klasifikasi

Sebelum *dataset* sentimen ini digunakan dalam pembuatan model, terlebih dahulu *dataset* di transformasi ke dalam bentuk matriks bilangan dengan menggunakan *Count Vectorizer*. Sehingga didapatkan kata beserta frekuensi dari setiap kata tersebut. Selanjutnya didapatkan distribusi data berdasarkan hasil frekuensi dari setiap kata yang ditunjukkan oleh Gambar 6. Pada Gambar 6 diketahui distribusi data menunjukkan nilai Mean sebesar 891, Q1 dan Q2 6897, dan Q3 1818.



Gambar 6. Distribusi Data

Tahap selanjutnya adalah membagi *dataset* menjadi dua bagian dengan proporsi 80% untuk data latih dan 20% untuk data uji. Kemudian, data latih tersebut di validasi menggunakan teknik *10-fold cross-validation* untuk mengukur kemampuan generalisasi dari model sehingga diperoleh 10 model yang berbeda. Selanjutnya dalam penelitian ini menggunakan beberapa algoritma klasifikasi, yaitu *Logistic Regression*, *Multinomial Naïve Bayes*, *Decision Tree*, *K-Nearest Neighbors*, *Random Forest*, dan *Support Vector Machine*. Terakhir model di evaluasi menggunakan 4 metrik evaluasi, yaitu *Accuracy*, *Precision*, *Recall* dan *F1-Score*. Hasil dari pengembangan model ditunjukkan dalam Tabel 4.

Tabel 4. Hasil Pengembangan Model

Algoritma	Accuracy	Precision	Recall	F1-Score
Logistic Regression	67%	67%	67%	67%
Multinomial Naive Bayes	66%	66%	66%	66%
Decision Tree	61%	61%	61%	61%
K-Nearest Neighbors	54%	54%	54%	54%
Random Forest	67%	67%	67%	67%
Support Vector Machine	69%	69%	69%	69%

Berdasarkan hasil model dengan beberapa algoritma, terlihat bahwa algoritma *Support Vector Machine* memiliki hasil terbaik dengan *accuracy* sebesar 69%. Namun, hasil tersebut di masih di bawah 70% sehingga diputuskan untuk menggunakan teknik SMOTE untuk meningkatkan kinerja model. *Synthetic Minority Oversampling Technique* (SMOTE) adalah teknik untuk menyeimbangkan kelas label dalam *dataset*. Hasil dari pengembangan model menggunakan teknik SMOTE ditunjukkan dalam Tabel 5.

Tabel 5. Hasil Pengembangan Model Menggunakan SMOTE

Algoritma	Accuracy	Precision	Recall	F1-Score
Logistic Regression	77%	77%	77%	77%
Multinomial Naive Bayes	72%	72%	72%	72%
Decision Tree	71%	71%	71%	71%
K-Nearest Neighbors	70%	70%	70%	70%
Random Forest	80%	80%	80%	80%
Support Vector Machine	82%	82%	82%	82%

Hasil dari penggunaan SMOTE mendapatkan peningkatan kinerja yang cukup signifikan di semua algoritma. *Logistic Regression* meningkat 10% dari 67% menjadi 77%. *Multinomial Naïve Bayes* meningkat 6% dari 66% menjadi 72%. *Decision Tree* meningkat 10% dari 61% menjadi 71%. *K-Nearest Neighbors* meningkat 16% dari 54% menjadi 70%. *Random Forest* meningkat 13% dari 67% menjadi 80%. Kemudian kinerja terbaik tetap di pegang oleh *Support Vector Machine* dengan peningkatan sebanyak 13% dari 69% menjadi 82%.

4. Kesimpulan

Berdasarkan penelitian yang telah dilakukan, hasil sentimen pada data tweet sangat berkaitan dengan kata kunci, isi dan jumlah tweet yang diambil. Sentimen netral menjadi permasalahan karena memiliki jumlah terbanyak dan isi tweet yang tidak relevan dan bersifat promosi. Kemudian, terdapat banyaknya tweet yang memiliki sentimen positif dibanding negatif, menunjukkan bahwa masyarakat cenderung menyukai dan ingin memiliki merek tas desainer. Tweet dengan sentimen negatif memiliki konteks yang menyatakan rasa tidak suka, benci, dan sarkasme terhadap suatu

merek, seperti rasa tidak suka terhadap harga dari masing-masing merek yang sangat tinggi.

Beberapa algoritma digunakan untuk pengembangan model pembelajaran mesin dalam penelitian ini. Support Vector Machine menjadi algoritma terbaik dibanding algoritma lainnya dengan accuracy 69%. Pada penelitian ini, telah dibuktikan bahwa teknik SMOTE dapat meningkatkan kinerja dari setiap model hingga 15%, di mana pada model dari algoritma Support Vector Machine meningkat sebesar 13%, dengan accuracy 82%.

Saran untuk penelitian lebih lanjut yang dapat dilakukan adalah dengan menentukan kata kunci pencarian yang lebih baik dari penelitian ini. Kemudian dapat dengan menambahkan data yang lebih banyak dari penelitian ini. Selain itu, dalam pengembangan model machine learning dapat menggunakan algoritma klasifikasi lain dan juga teknik-teknik lain yang bertujuan untuk meningkatkan performa model.

Daftar Rujukan

- [1] Hidapenta, D., & Dewi, D. A. (2021). Peran Pkn Mengatasi Fenomena Kecintaan Produk Luar Yang Terjadi Di Indonesia. *Jurnal Kewarganegaraan*, 5(1), 168–175. <https://doi.org/10.31316/jk.v5i1.1401>
- [2] Yanuarsari, D. H. (2015). Analisis Minat Beli Wanita Terhadap Produk Tas Bermerek Original di Tengah Komoditi Produksi Tas Bermerek Tiruan Produksi Produsen Lokal. *ANDHARUPA: Jurnal Desain Komunikasi Visual & Multimedia*, 1(02), 110–121. <https://doi.org/10.33633/andharupa.v1i02.961>
- [3] Vijay, S., Prashar, S., & Gupta, S. (2019). An examination of the role of review valence and review source in varying consumption contexts on purchase decision. *Journal of Retailing and Consumer Services*, 52(January), 101734. <https://doi.org/10.1016/j.jretconser.2019.01.003>
- [4] Chetoui, Y., Benlafqih, H., & Lebdaoui, H. (2020). *How fashion influencers contribute Influencers to consumers' purchase intention*. 24(3), 361–380. <https://doi.org/10.1108/JFMM-08-2019-0157>
- [5] Yaacob, A., Gan, J. L., & Yusuf, S. (2021). The Role of Online Consumer Review, Social Media Advertisement and Influencer Endorsement on Purchase Intention of Fashion Apparel during COVID-19. *Journal of Content, Community & Communication*, 14(8), 17–33. <https://doi.org/10.31620/JCCC.12.21/03>
- [6] Lee, M., Cai, Y. (Maggie), DeFranco, A., & Lee, J. (2020). Exploring influential factors affecting guest satisfaction: Big data and business analytics in consumer-generated reviews. *Journal of Hospitality and Tourism Technology*, 11(1), 137–153. <https://doi.org/10.1108/JHTT-07-2018-0054>
- [7] Weistein, F. L., Song, L., Andersen, P., & Zhu, Y. (2017). Examining impacts of negative reviews and purchase goals on consumer purchase decision. *Journal of Retailing and Consumer Services*, 39(June), 201–207. <https://doi.org/10.1016/j.jretconser.2017.08.015>
- [8] Diekson, Z. A., Prakoso, M. R. B., Putra, M. S. Q., Syaputra, M. S. A. F., Achmad, S., & Sutoyo, R. (2023). Sentiment analysis for customer review: Case study of Traveloka. *Procedia Computer Science*, 216(2022), 682–690. <https://doi.org/10.1016/j.procs.2022.12.184>
- [9] Prananda, A. R., & Thalib, I. (2020). Sentiment Analysis for Customer Review: Case Study of GO-JEK Expansion. *Journal of Information Systems Engineering and Business Intelligence*, 6(1), 1. <https://doi.org/10.20473/jisebi.6.1.1-8>
- [10] Rambocas, M., & Pacheco, B. G. (2018). Online sentiment analysis in marketing research: a review. *Journal of Research in Interactive Marketing*, 12(2), 146–163. <https://doi.org/10.1108/JRIM-05-2017-0030>
- [11] Lee, D., Jo, J.-C., & Lim, H.-S. (2017). User Sentiment Analysis on Amazon Fashion Product Review Using Word Embedding. *Journal of the Korea Convergence Society*, 8(4), 1–8. <https://doi.org/10.15207/jkcs.2017.8.4.001>
- [12] Pantano, E., Giglio, S., & Dennis, C. (2019). Making sense of consumers' tweets: Sentiment outcomes for fast fashion retailers through Big Data analytics. *International Journal of Retail and Distribution Management*, 47(9), 915–927. <https://doi.org/10.1108/IJRDM-07-2018-0127>
- [13] Susilayasa, I. M. A., Eka Karyawati, A. A. I., Astuti, L. G., Rahning Putri, L. A. A., Arta Wibawa, I. G., & Ari Mogi, I. K. (2022). Analisis Sentimen Ulasan E-Commerce Pakaian Berdasarkan Kategori dengan Algoritma Convolutional Neural Network. *JELIKU (Jurnal Elektronik Ilmu Komputer Udayana)*, 11(1), 1. <https://doi.org/10.24843/jlk.2022.v11i01.p01>
- [14] Muslim, M. A., Dasril, Y., Alamsyah, A., & Mustaqim, T. (2021). Bank predictions for prospective long-term deposit investors using machine learning LightGBM and SMOTE. *Journal of Physics: Conference Series*, 1918(4). <https://doi.org/10.1088/1742-6596/1918/4/042143>
- [15] Peffers, K., Tuunanen, T., Gengler, C. E., Rossi, M., Hui, W., Virtanen, V., & Bragge, J. (2020). Design Science Research Process: A Model for Producing and Presenting Information Systems Research. *ArXiv, abs/2006.0*.
- [16] Huda, M. N. (2023). *Bag Brands Sentiment Dataset [Data set]*. Zenodo. <https://doi.org/10.5281/zenodo.7679325>
- [17] Lourdasamy, R., & Abraham, S. (2018). A Survey on Text Pre-processing Techniques and Tools. *International Journal of Computer Sciences and Engineering*, 06(03), 148–157. <https://doi.org/10.26438/ijcse/v6si3.148157>
- [18] Turki, T., & Roy, S. S. (2022). Novel Hate Speech Detection Using Word Cloud Visualization and Ensemble Learning Coupled with Count Vectorizer. *Applied Sciences (Switzerland)*, 12(13). <https://doi.org/10.3390/app12136611>
- [19] Pamungkas, M. R. S. P., Huda, M. N., Fauzan, D. A., Itsna, A. H., & Al-Hijri, F. M. (2022). Sistem Klasifikasi Otomatis Dengan Konsep Machine Learning As A Service (MLaaS) Pada Kasus Pesan Berindikasi Cyberbullying. *ILKOMNIKA: Journal of Computer and Applied Informatics*, 4(3), 252–261. <https://doi.org/10.28926/ilkomnika.v4i3.522>
- [20] Chawla, N. V., Bowyer, K. W., Hall, L. O., Department, & Kegelmeyer, W. P. (2002). SMOTE: Synthetic Minority Over-sampling Technique. *Journal of Artificial Intelligence Research (JAIR)*, 16, 321–357. <https://doi.org/10.1613/jair.953>
- [21] Alkharabsheh, K., Alawadi, S., Kebande, V. R., Crespo, Y., Fernández-Delgado, M., & Taboada, J. A. (2022). A comparison of machine learning algorithms on design smell detection using balanced and imbalanced dataset: A study of God class. *Information and Software Technology*, 143, 106736. <https://doi.org/10.1016/j.infsof.2021.106736>
- [22] Hemmatian, F., & Sohrabi, M. K. (2017). A survey on classification techniques for opinion mining and sentiment analysis. *Artificial Intelligence Review*, 52(3), 1495–1545. <https://doi.org/10.1007/s10462-017-9599-6>
- [23] Kabir, A. I., Ahmed, K., & Karim, R. (2020). Word Cloud and Sentiment Analysis of Amazon Earphones Reviews with R Programming Language. *Informatica Economica*, 24(4/2020), 55–71. <https://doi.org/10.24818/issn14531305/24.4.2020.05>