

ANALISIS SENTIMEN ULASAN PELANGGAN STARBUCKS MENGUNAKAN NAÏVE BAYES DAN SUPPORT VECTOR MACHINE

Fredy Liestianto

Program Studi Teknik Informatika, Fakultas Teknologi Informasi, Universitas Tarumanagara
Jl. Letjen S.Parman No.1, Jakarta Barat, DKI Jakarta, Indonesia 11410
e-mail: fredyliestianto@gmail.com

ABSTRAK

Penelitian ini menerapkan pendekatan klasifikasi Naïve Bayes dan Support Vector Machine (SVM) untuk menganalisis sentimen ulasan pelanggan Starbucks. Data ulasan pelanggan diekstrak dari platform online, kemudian diberikan tahap preprocessing. Naïve Bayes dan SVM digunakan untuk mengklasifikasikan sentimen menjadi positif, negatif, atau netral, dengan evaluasi performa menggunakan metrik klasifikasi. Hasil menunjukkan bahwa kombinasi keduanya memberikan analisis sentimen yang lebih akurat, berpotensi membantu perusahaan memahami persepsi pelanggan dan meningkatkan kualitas produk serta layanan. Penelitian ini berkontribusi pada pengembangan metode analisis sentimen yang efektif dalam konteks bisnis.

Kata kunci: Analisis Sentimen, Ulasan Pelanggan, Starbucks, Naïve Bayes, Support Vector Machine.

ABSTRACT

This research applies the classification approach of Naive Bayes and Support Vector Machine (SVM) to analyze customer sentiment in Starbucks reviews. Customer review data is extracted from online platforms, followed by a preprocessing stage. Naive Bayes and SVM are employed to classify sentiments into positive, negative, or neutral, with performance evaluation using classification metrics. The results indicate that their combination provides more accurate sentiment analysis, potentially aiding companies in understanding customer perceptions and enhancing product and service quality. This research contributes to the development of effective sentiment analysis methods in a business context.

Keywords: Sentiment Analysis, Customer Reviews, Starbucks, Naïve Bayes, Support Vector Machine.

1 PENDAHULUAN

Dalam era digital, ulasan pelanggan secara daring menjadi sumber informasi yang penting bagi perusahaan dalam memahami persepsi konsumen terhadap produk dan layanan mereka. Starbucks, sebagai salah satu merek kopi terkemuka dengan kehadiran global, tidak terkecuali dari fenomena ini. Oleh karena itu, penelitian ini bertujuan untuk menerapkan pendekatan klasifikasi menggunakan algoritma Naive Bayes dan Support Vector Machine (SVM) guna menganalisis sentimen yang terkandung dalam ulasan pelanggan Starbucks.

Data ulasan pelanggan yang diperoleh dari platform daring merupakan sumber informasi yang berharga untuk memahami kepuasan pelanggan, preferensi konsumen, serta potensi perbaikan yang dapat dilakukan oleh perusahaan. Proses analisis sentimen ini melibatkan penerapan algoritma Naive Bayes dan SVM sebagai metode klasifikasi, yang diterapkan pada tahap preprocessing data guna meningkatkan akurasi hasil analisis.

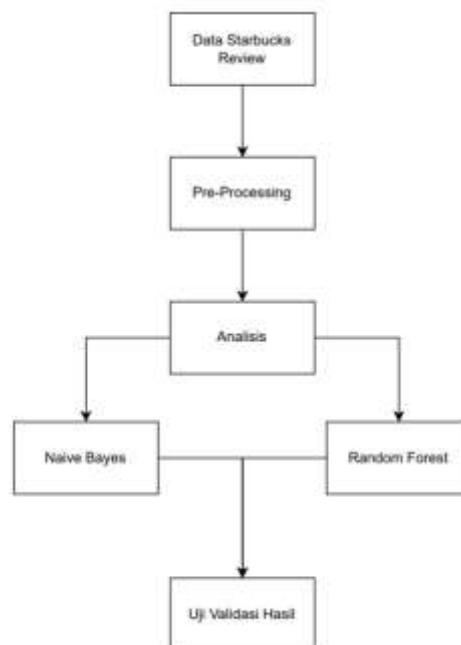
Penelitian ini diharapkan dapat memberikan kontribusi dalam pemahaman lebih mendalam tentang dinamika sentimen pelanggan terhadap Starbucks, sekaligus menghadirkan metodologi yang efektif dalam menganalisis ulasan pelanggan secara menyeluruh. Hasil dari penelitian ini diharapkan dapat menjadi landasan bagi perusahaan untuk meningkatkan pengalaman pelanggan, serta mengoptimalkan kualitas produk dan layanan mereka. Dengan demikian, penelitian ini menjadi relevan dalam konteks pengembangan strategi bisnis yang berfokus pada kebutuhan dan harapan konsumen.

2 TINJAUAN LITERATUR

Proses klasifikasi sentimen terhadap kepuasan pelanggan Starbucks menghasilkan kategori netral. Dari analisis ulasan yang menggunakan kata kunci "starbuck" atau "starbucks" atau "#starbucks", diperoleh hasil sebagai berikut: komentar positif terdapat sebanyak 476 tweet dengan persentase sebesar 19,2%, komentar netral sebanyak 1743 tweet dengan persentase sebesar 70,3%, dan komentar negatif sebanyak 258 tweet dengan persentase sebesar 10,4%. Dengan demikian, berdasarkan perhitungan polarity tersebut, kesimpulannya adalah ulasan komentar terhadap Starbucks dapat dikategorikan sebagai puas.[1]

Dari hasil studi yang dilakukan menggunakan 265 data komentar pada Google Play Store, dengan menerapkan tiga metode klasifikasi, yaitu k-Nearest Neighbor (K-NN), Support Vector Machine (SVM), dan Random Forest, diperoleh hasil sebagai berikut: Metode klasifikasi K-NN menunjukkan nilai akurasi sebesar 89,0%, presisi sebesar 89,7%, dan recall sebesar 87,5%. Metode klasifikasi Random Forest menunjukkan nilai akurasi sebesar 83,0%, presisi sebesar 85,7%, dan recall sebesar 81,4%. Metode klasifikasi SVM menunjukkan nilai akurasi sebesar 89,4%, presisi sebesar 89,5%, dan recall sebesar 89,7%. Dengan demikian, berdasarkan hasil tersebut, dapat disimpulkan bahwa metode klasifikasi terbaik pada studi ini adalah SVM, dengan nilai akurasi 89,4%, presisi 89,5%, dan recall 89,7%.[2]

3 METODE PENELITIAN



Gambar 1. Langkah – Langkah Klasifikasi Teks

Tentu saja, dalam penelitian ini, langkah pertama yang harus dilakukan adalah mengumpulkan data. Setelah data ulasan berhasil dikumpulkan dari Kaggle, langkah-langkah selanjutnya diambil agar hasil yang diperoleh optimal. Data tersebut kemudian diolah melalui proses pembersihan (*cleansing*) dan pemberian label. Selanjutnya, masuk ke tahap pre-processing yang mencakup *case folding*, *cleaning*, *tokenizing*, *stopword removal*, *stemming*, *labeling*, dan berlanjut ke tahap analisis. Proses analisis melibatkan penggunaan berbagai alat dan metode, termasuk training data, testing data, serta penerapan tiga metode klasifikasi, yaitu Naïve Bayes dan Support Vector Machine. Gambar 1 memberikan gambaran visual tentang langkah-langkah dalam analisis klasifikasi sentimen yang diusulkan dalam penelitian ini.[2]

3.1 *Confusion Matrix*

Penggunaan *Confusion Matrix* adalah salah satu metode untuk mengukur kinerja pada suatu metode pengelompokkan. *Confusion Matrix* bekerja dengan membandingkan dataset yang telah diuji ke dalam empat kategori. Berikut adalah rumus untuk melakukan pengukuran nilai akurasi pada klasifikasi teks:[1]

$$\text{Akurasi} = \frac{TP+TN}{Total}$$

$$\text{Presisi} = \frac{TP}{TP+FP}$$

$$\text{Recall} = \frac{TP}{TP+FN}$$

3.2 *Naïve Bayes*

Algoritma *Naïve Bayes* merupakan salah satu algoritma yang umum digunakan dalam penelitian untuk mengatasi klasifikasi pada teks dengan memanfaatkan perhitungan probabilitas. Naive Bayes menjadi pilihan yang populer terutama dalam situasi di mana terdapat jumlah data yang besar dan keberadaan data yang kosong. Algoritma ini menggunakan perhitungan probabilitas dan statistik untuk melakukan prediksi terkait suatu data, khususnya dalam pengujian yang telah dilakukan sebelumnya. Konsep ini sering disebut sebagai missing value, di mana algoritma Naive Bayes mampu menangani data yang kurang atau tidak lengkap.[1]

3.3 *Support Vector Machine*

Support Vector Machine (SVM) adalah metode pembelajaran terawasi yang mengamati data dan mengidentifikasi pola untuk pengelompokan. SVM dianggap sebagai metode yang sangat efektif karena mampu memproses data dengan efisien. Meskipun demikian, SVM memiliki kekurangan dalam menentukan parameter atau fitur yang optimal. Untuk mengatasi ini, SVM diusulkan untuk mengaplikasikan klasifikasi sentiment menggunakan model mesin klasifikasi pembelajaran domain dengan memanfaatkan berbagai fitur tekstual, termasuk data. Proses tersebut melibatkan tiga skema penimbangan yang berbeda untuk memahami dampak penimbangan terhadap akurasi pengklasifikasi. Studi masa depan bertujuan untuk memberikan pemahaman tambahan dan hasil yang dapat meningkatkan kinerja pengklasifikasi.[2]

3.4 *Jupyter Notebook*

Jupyter adalah aplikasi web gratis yang digunakan untuk membuat dan berbagi dokumen yang menggabungkan kode, hasil perhitungan, visualisasi, dan teks. Platform ini memberikan kemudahan bagi individu dan tim untuk melakukan analisis data secara interaktif dan eksploratif. Melalui Jupyter, pengguna dapat menciptakan dokumen yang mengintegrasikan elemen-elemen tersebut untuk menjelaskan data, menghasilkan wawasan yang signifikan, dan berkomunikasi secara efisien mengenai hasil analisis. Jupyter, singkatan dari tiga bahasa pemrograman utama, yaitu Julia (Ju), Python (Py), dan R, memiliki peran kunci dalam ekosistem analisis data dan ilmu data. Kombinasi ini memungkinkan data scientist untuk mengakses beragam alat dan pustaka yang diperlukan untuk menjalankan analisis data yang canggih. Peran utama Jupyter adalah membantu pengguna dalam membuat narasi komputasi yang menjelaskan arti data dalam dokumen mereka dan memberikan wawasan mendalam mengenai data tersebut. Dengan demikian, Jupyter menjadi alat yang efektif untuk menjelaskan temuan dalam analisis data, membuat laporan, atau bahkan mengajar konsep ilmu data kepada orang lain.[6]

4 HASIL DAN PEMBAHASAN

4.1 Dataset

Sebuah dataset merujuk pada sekelompok data yang terdiri dari beragam objek, beserta atribut-atribut yang merepresentasikan karakteristik dari setiap objek tersebut. Dataset ini dapat terdiri dari berbagai file atau dokumen yang mengandung informasi yang dapat dianalisis dan dieksplorasi. Dalam bidang statistik, dataset umumnya berasal dari pengamatan yang dilakukan secara aktual dengan cara mengambil sampel dari populasi yang lebih besar.[5] Dataset yang digunakan dalam konteks ini berasal dari website Kaggle dengan nama file starbucks.csv, yang mengandung 813 data. Sebagai bagian dari penggunaan dataset tersebut, sekitar 30% dari total data, yaitu sekitar 244 data, akan diambil sebagai data uji untuk keperluan analisis atau pemodelan selanjutnya.

4.2 Preprocessing data

Data-data yang diambil dari website Kaggle akan diolah kembali menggunakan metode pre-processing. Metode *pre-processing* bertujuan untuk melakukan eliminasi data yang tidak sesuai dengan parameter yang telah ditentukan. Berikut adalah tahapan *pre-processing* yang terdiri dari:[1]

4.2.1 Case Folding

Case folding adalah proses mengubah semua huruf menjadi huruf kecil. Hal ini dilakukan karena data yang dimiliki tidak selalu konsisten dan terstruktur dalam menggunakan huruf kapital. Berikut adalah data yang dihasilkan dari proses case folding:[1]

Review	
Amber and LaDonna at the Starbucks on Southwes...	
** at the Starbucks by the fire station on 436...	
I just wanted to go out of my way to recognize...	
Me and my friend were at Starbucks and my card...	
I'm on this kick of drinking 5 cups of warm wa...	

After Case Folding:	
0	amber and ladonna at the starbucks on southwes...
1	** at the starbucks by the fire station on 436...
2	i just wanted to go out of my way to recognize...
3	me and my friend were at starbucks and my card...
4	i'm on this kick of drinking 5 cups of warm wa...
Name: case_folded_text, dtype: object	

Gambar 2. Case Folding

Gambar 2 menunjukkan bahwa telah dilakukan perubahan pada semua kata yang terdapat pada teks kalimat dari huruf "a" sampai "z" menjadi huruf kecil.[1]

4.2.2 Cleaning

Cleaning adalah proses menyiapkan data dengan cara menghapus atau memodifikasi data yang salah, tidak relevan, tidak akurat, duplikat, maupun yang tidak terformat. Setelah melewati tahap ini, barulah data siap untuk diolah. Berikut adalah data yang dihasilkan dari proses cleaning:[1]

After Case Folding:	
0	amber and ladonna at the starbucks on southwes...
1	** at the starbucks by the fire station on 436 i...
2	i just wanted to go out of my way to recognize...
3	me and my friend were at starbucks and my card...
4	i'm on this kick of drinking 5 cups of warm wa...
Name: case_folded_text, dtype: object	

After Cleaning:	
0	amber and ladonna at the starbucks on southwes...
1	at the starbucks by the fire station on 436 i...
2	i just wanted to go out of my way to recognize...
3	me and my friend were at starbucks and my card...
4	im on this kick of drinking 5 cups of warm wat...
Name: cleaned_text, dtype: object	

Gambar 3. Cleaning

Gambar 3 menunjukkan bahwa dilakukan penghapusan kata pada teks yang tidak memiliki arti, seperti kata penghubung, tanda baca, simbol atau karakter, angka, emoji, dan URL.[1]

4.2.3 Tokenizing

Tokenizing adalah metode untuk melakukan pemisahan kata dalam suatu kalimat dengan tujuan untuk proses analisis teks lebih lanjut. Berikut adalah data yang dihasilkan dari proses *tokenizing*: [1]

```
After Cleaning:
0  amber and ladonna at the starbucks on southwest...
1  at the starbucks by the fire station on 436 i...
2  i just wanted to go out of my way to recognize...
3  me and my friend were at starbucks and my card...
4  im on this kick of drinking 5 cups of warm wat...
Name: cleaned_text, dtype: object

After Tokenizing:
0  [amber, and, ladonna, at, the, starbucks, on, ...
1  [at, the, starbucks, by, the, fire, station, o...
2  [i, just, wanted, to, go, out, of, my, way, to...
3  [me, and, my, friend, were, at, starbucks, and...
4  [im, on, this, kick, of, drinking, 5, cups, of...
Name: tokenized_text, dtype: object
```

Gambar 4. Tokenizing

Gambar 4 menunjukkan bahwa telah dilakukan pemisahan teks pada kalimat menjadi satu kata sesuai dengan penulisan spasi pada teks.[1]

4.2.4 Stopword Removal

Stopword Removal adalah proses menghilangkan kata penghubung. Tujuannya adalah mengurangi jumlah kata dalam sebuah dokumen untuk meningkatkan kecepatan dan performa. Berikut adalah data yang dihasilkan dari proses *Stopword Removal*: [1]

```
After Tokenizing:
0  [amber, and, ladonna, at, the, starbucks, on, ...
1  [at, the, starbucks, by, the, fire, station, o...
2  [i, just, wanted, to, go, out, of, my, way, to...
3  [me, and, my, friend, were, at, starbucks, and...
4  [im, on, this, kick, of, drinking, 5, cups, of...
Name: tokenized_text, dtype: object

After Stopword Removal:
0  [amber, ladonna, starbucks, southwest, parkway...
1  [starbucks, fire, station, 436, altamonte, spr...
2  [wanted, go, way, recognize, starbucks, employ...
3  [friend, starbucks, card, didnt, work, thankfu...
4  [im, kick, drinking, 5, cups, warm, water, wor...
Name: no_stopword_text, dtype: object
```

Gambar 5. Stopword Removal

Gambar 5 menggambarkan tahap dalam memilih kata-kata dari teks yang tidak memiliki arti, seperti kata-kata keterangan, kata sambung, dan lainnya yang tidak diperlukan dalam proses pemodelan data.[1]

4.2.5 Stemming

Stemming adalah proses menghilangkan imbuhan pada suatu kata. Berikut adalah data yang dihasilkan dari proses *stemming*: [1]

```
After Stopword Removal:
0  [amber, ladonna, starbucks, southwest, parkway...
1  [starbucks, fire, station, 436, altamonte, spr...
2  [wanted, go, way, recognize, starbucks, employ...
3  [friend, starbucks, card, didnt, work, thankfu...
4  [im, kick, drinking, 5, cups, warm, water, wor...
Name: no_stopword_text, dtype: object

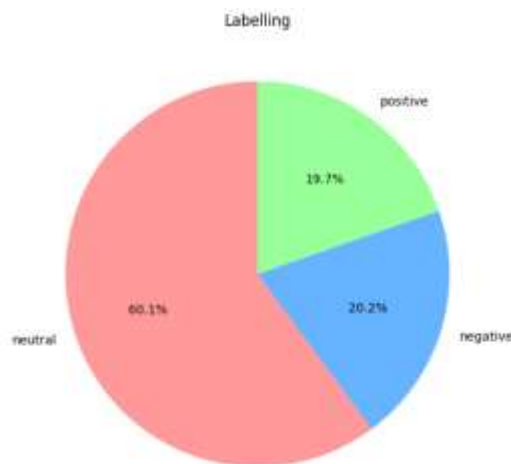
After Stemming:
0  [amber, ladonna, starbuck, southwest, parkway...
1  [starbuck, fire, station, 436, altamont, sprin...
2  [want, go, way, recogn, starbuck, employe, bil...
3  [friend, starbuck, card, didnt, work, thank, w...
4  [im, kick, drink, 5, cup, warm, water, work, i...
Name: stemmed_text, dtype: object
```

Gambar 6. Stemming

Gambar 6 menunjukkan bahwa data yang dihasilkan dari proses penghilangan stopwords akan menghilangkan kata-kata tambahan atau imbuhan.[1]

4.2.6 Labelling

Dalam langkah ini, data yang dihasilkan dari proses Stopword Removal akan dihitung berdasarkan polaritas ulasan pelanggan yang diperoleh. Hal ini akan menghasilkan tiga kategori pada klasifikasi, yaitu: label negatif untuk ulasan dengan nilai kurang dari 3, label netral untuk ulasan dengan nilai sama dengan 3, dan label positif untuk ulasan dengan nilai lebih dari 3.[1]



Gambar 7. Labelling

Diagram lingkaran (*Pie chart*) pada Gambar 7 memberikan gambaran tentang klasifikasi ulasan pelanggan terhadap Starbucks yang diperoleh dari Kaggle. Dapat dilihat bahwa label-label klasifikasi memiliki persentase sebagai berikut: label positif mencapai 19,7%, label negatif mencapai 20,2%, dan label netral mencapai 60,1%.[1]

4.3. Evaluasi

4.3.1 Hasil dari Naïve Bayes

Confusion Matrix Naive Bayes:

	Predicted Negative	Predicted Neutral	Predicted Positive
Actual Negative	0	44	0
Actual Neutral	0	151	0
Actual Positive	0	49	0

Gambar 8. Confusion Matrix Naïve Bayes

Gambar 8 memberikan penjelasan mengenai hasil *Confusion Matrix* untuk setiap kelas dalam kondisi kelas sebenarnya (*actual*), di mana model dapat memprediksi dengan benar. Penjelasan ini dapat diuraikan sebagai berikut: tidak ada data yang diklasifikasikan sebagai *True Positive* (TP) untuk kelas negatif, terdapat 151 data yang diidentifikasi sebagai TP untuk kelas netral, dan tidak ada data yang terklasifikasikan sebagai TP untuk kelas positif. Evaluasi pada matriks dilakukan untuk menghitung nilai akurasi menggunakan algoritma *Naive Bayes* dengan menggunakan bahasa pemrograman *Python* di *Jupyter Notebook*, dan rincian perhitungan tersebut dapat ditemukan pada gambar berikut: [1]

Naive Bayes Classification Report:				
	precision	recall	f1-score	support
negative	0.00	0.00	0.00	44
neutral	0.62	1.00	0.76	151
positive	0.00	0.00	0.00	49
accuracy			0.62	244
macro avg	0.21	0.33	0.25	244
weighted avg	0.38	0.62	0.47	244

Gambar 9. Laporan Klasifikasi Naïve Bayes

Gambar 9 menunjukkan hasil perhitungan matriks menggunakan bahasa pemrograman *Python* di *Jupyter Notebook*, yaitu sebesar 0,62 atau 62%. Di bawah ini disajikan penjelasan perhitungan nilai akurasi secara manual berdasarkan perhitungan matriks tersebut: [1]

$$\text{Akurasi} = \frac{151}{244} \times 100 = 62\%$$

4.3.2 Hasil dari *Support Vector Machine*

Confusion Matrix Support Vector Machine:

	Predicted Negative	Predicted Neutral	Predicted Positive
Actual Negative	0	44	0
Actual Neutral	0	151	0
Actual Positive	0	44	5

Gambar 10. Confusion Matrix Support Vector Machine

Pada Gambar 10, terdapat penjelasan mengenai hasil *Confusion Matrix* dari setiap kelas dalam kondisi kelas sebenarnya (*actual*), di mana model mampu membuat prediksi yang tepat. Penjelasan ini dapat dirinci sebagai berikut: tidak ada data yang diklasifikasikan sebagai *True Positive* (TP) untuk kelas negatif, terdapat 151 data yang diidentifikasi sebagai TP untuk kelas neutral, dan ada 5 data yang terklasifikasikan sebagai TP untuk kelas positif. Evaluasi pada matriks dilakukan untuk menghitung nilai akurasi menggunakan algoritma *Naive Bayes* dengan menggunakan bahasa pemrograman *Python* di *Jupyter Notebook*, dan rincian perhitungan tersebut dapat ditemukan pada gambar berikut: [1]

Support Vector Machine Classification Report:

	precision	recall	f1-score	support
negative	0.00	0.00	0.00	44
neutral	0.63	1.00	0.77	151
positive	1.00	0.10	0.19	49
accuracy			0.64	244
macro avg	0.54	0.37	0.32	244
weighted avg	0.59	0.64	0.52	244

Gambar 11. Laporan Klasifikasi Support Vector Machine

Dalam Gambar 11, tampak hasil perhitungan matriks menggunakan bahasa pemrograman *Python* pada *Jupyter Notebook*, yakni 0,64 atau setara dengan 64%. Berikut adalah penjelasan manual perhitungan nilai akurasi berdasarkan perhitungan matriks tersebut [1]

$$\text{Akurasi} = \frac{156}{244} \times 100 = 64\%$$

5 KESIMPULAN

Hasil dari penelitian ini menunjukkan bahwa melalui proses klasifikasi sentimen terhadap ulasan pelanggan Starbucks, ditemukan bahwa kategori netral mendominasi. Hasil tersebut mencakup komentar positif sebanyak 19,7%, komentar netral sebanyak 60,1%, dan komentar negatif sebanyak 20,2%. Dengan demikian, dapat disimpulkan bahwa berdasarkan perhitungan polaritas, umpan balik pelanggan terhadap Starbucks dikategorikan sebagai puas. Dari penelitian ini, dapat disimpulkan bahwa kinerja algoritma *Support Vector Machine* lebih unggul daripada algoritma *Naive Bayes*, seperti yang terlihat dari rincian berikut. Akurasi algoritma *Support Vector Machine* mencapai 64%, sedangkan algoritma *Naive Bayes* hanya mencapai 62%.

DAFTAR PUSTAKA

- [1] T. A. Q. Putri, A. Triayudi, R. T. Aldisa, "Implementasi Algoritma Decision Tree dan Naïve Bayes Untuk Klasifikasi Sentimen Terhadap Kepuasan Pelanggan Starbucks", *Journal of Information System Research (JOSH)*, vol. 4, no. 2, pp 641–649, 2023.
- [2] S. Watmah, Suryanto, Martias, "Komparasi Metode K-NN, Support Vector Machine, Dan Random Forest Pada E-Commerce Shopee", *INSANtek – Jurnal Inovasi dan Sains Teknik Elektro*, vol. 2, no. 1, 2021.
- [3] J. P. Putra, T. Janji, R. Sitinjak, "Pengaruh kualitas produk dan harga terhadap kepuasan pelanggan Kopi Starbucks di Summarecon Mall Kelapa Gading 3", *Jurnal Ilmiah Akuntansi dan Keuangan*, vol. 4, no. 8, p. 2022, 2022, [Online]. Available: <https://journal.ikopin.ac.id/index.php/fairvalue>.
- [4] A. S. D. Herlambang, E. Komara, A. S. Herlambang, "Pengaruh Kualitas Produk, Kualitas Pelayanan, Dan Kualitas Promosi Terhadap Kepuasan Pelanggan (Studi kasus pada Starbucks Coffee Reserve Plaza Senayan)," *Jurnal Ekonomi, Manajemen dan Perbankan (Journal of Economics, Management and Banking)*, vol. 7, no. 2, 2021.
- [5] D. A. A. Kurniawan, E. Utami, H. A. Fatta, "Analisis Sentimen Pada Opini Pengguna Aplikasi Qasir Menggunakan Support Vector Machine Dan Random Forest", *TEKNIMEDIA Teknologi Informasi dan Multimedia* 4, vol. 1, no. 1-8, 2023.
- [6] M. I. Medina, "Penting untuk Data Scientist, Ketahui Apa Saja Fungsi dan Fitur Jupyter", 2022. [Online]. Available: <https://qlints.com/id/lowongan/Jupyter-adalah/>.
- [7] B. Prasajo, E. Haryatmi, "Analisa Prediksi Kelayakan Pemberian Kredit Pinjaman dengan Metode Random Forest", *Jurnal Nasional Teknologi Dan Sistem Informasi*, vol. 7, no. 2, 2021.
- [8] Yunitasari, H. S. Hopipah, R. Mayasari, "Optimasi Backward Elimination untuk Klasifikasi Kepuasan Pelanggan Menggunakan Algoritme k-Nearest Neighbor (k-NN) dan Naïve Bayes", *Technomedia Journal (TMJ)*, vol. 6, no. 1, 2021, doi: 10.33050/tmj.v6i1.
- [9] A. K. Febrian, Y. H. Chrisnanto, D. Pupita, N. Sabrina, J. A. Yani, "Studi Komparasi Metode Klasifikasi K-Nearest Neighbor dan Naïve Bayes dalam Mengidentifikasi Kepuasan Pelanggan Terhadap Produk", *Seminar Nasional Teknik Elektro, Sistem Informasi, dan Teknik Informatika*, p. 333, 2022, doi: 10.31284/p.snestik.2022.2717.
- [10] M. A. Nurhakim, Y. Widiastih, N. Chamidah, "Analisis Sentimen Terhadap Ulasan Kepuasan Pelanggan Pada Marketplace Tokopedia Di Jejaring Sosial Twitter Menggunakan Algoritma Naïve Bayes", *Jurnal Ilmiah Betrik*, vol. 14, no. 2, 2022.
- [11] N. Saurina, T. Rahayuningsih, L. Retnawati, "Analisis Sentimen Ulasan Pelanggan Batik Ecoprint Menggunakan Naïve Bayes Dan KNN Classifier", *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 9, no. 2, 2022.
- [12] D. Sepri, "Penerapan Algoritma Naïve Bayes Untuk Analisis Kepuasan Penggunaan Aplikasi Bank", *Journal of Computer System and Informatics (JoSYC)*, vol. 2, no. 1, 2020.
- [13] M. I. Petiwi, A. Triayudi, I. D. Sholihati, "Analisis Sentimen Gofood Berdasarkan Twitter Menggunakan Metode Naïve Bayes dan Support Vector Machine", *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol.6, no.1, 2022.
- [14] D. J. Pangestu, A. Kodar, "Implementasi Multinomial Naïve Bayes Untuk Klasifikasi Sentimen Terhadap Pelayanan Perusahaan Otobus Menggunakan Data Facebook (Studi Kasus: Grup Facebook Murni Jaya Lovers)", *Jurnal Informatika: Jurnal pengembangan IT (JPIT)*, vol.7, no.3, 2022.
- [15] A. Syafii, G. Dwilestari, A. Ajiz, "Komparasi Algoritma Naïve Bayes Dan Algoritma C4.5 Dalam Klasifikasi Pelanggan Produk Indihome", *JURSIMA Jurnal Sistem Informasi dan Manajemen*, vol. 10, no. 2, 2022.
- [16] W. I. Rahayu, A. Anindita, M. N. Fauzan, "Penentuan Validasi Data Pemilih Dan Klasifikasi Hasil Pemilu Dprd Kab.Bone Untuk Memprediksi Partai Pemenang Menggunakan Metode Naive Bayes", *Jurnal Teknik Informatika*, vol. 14, no. 1, 2022.
- [17] A. Tangkelayuk, E. Mailoa, "Klasifikasi Kualitas Air Menggunakan Metode KNN, Naïve Bayes Dan Decision Tree", *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 9, no. 2, 2022.
- [18] Farmawati, Narti, "Perbandingan Algoritma C4.5 dan Naive Bayes dalam Klasifikasi Tingkat Kepuasan Mahasiswa Terhadap Pembelajaran Daring", *JTIM : Jurnal Teknologi Informasi dan Multimedia*, vol. 4, no. 1, 2022.
- [19] M. A. Djamaludin, A. Triayudi, E. Mardiani, "Analisis Sentimen Tweet KRI Nanggala 402 di Twitter menggunakan Metode Naïve Bayes Classifier", *Jurnal JTIK (Jurnal Teknologi Informasi dan Komunikasi)*, vol. 6, no. 2, 2022.
- [20] A. Rozaqi, A. Triayudi, R. T. Aldisa, "Analisis Sentimen Vaksinasi Booster Berdasarkan Twitter Menggunakan Algoritma Naïve Bayes dan K-NN", *Jurnal Sistem Komputer dan Informatika (JSON)*, vol. 4, no. 1, 2022.