

ANALISIS SENTIMEN DENGAN ALGORITMA NAÏVE BAYES CLASSIFIER TERHADAP ULASAN APLIKASI MY F&B ID

Nabila Ayu Puspita, Fetty Tri Anggraeny, Agung Mustika Rizki

Teknik Informatika, Universitas Pembangunan Nasional Veteran Jawa Timur

Jl. Rungkut Madya No.1, Gn. Anyar, Kec. Gn. Anyar, Surabaya, Jawa Timur

19081010006@student.upnjatim.ac.id

ABSTRAK

Di era teknologi yang semakin canggih mendorong semua aspek bidang kehidupan ikut dalam perkembangannya, salah satunya pada bidang makanan atau kuliner. Kecanggihan hal tersebut ialah membuat perusahaan makanan berlomba-lomba untuk membuat aplikasi *mobile* untuk memberi kemudahan pada masyarakat terhadap pembelian produk yang mereka jual. Salah satu bentuk implementasi nyata dari hal tersebut dengan adanya Aplikasi My F&B ID. Pada *Google Play Store* selain dapat mengunduh aplikasi, masyarakat juga dapat memberi pendapat terhadap aplikasi tersebut. Pendapat yang berupa memberi ulasan masukkan, kritik, pujian, keluhan terhadap aplikasi yang dapat meningkatkan kualitas dan pelayanan yang ditawarkan aplikasi. Maka dalam penelitian ini akan dilakukan analisis sentimen terhadap ulasan aplikasi My F&B ID pada *Google Playstore* menggunakan algoritma Naïve Bayes Classifier untuk mengetahui permasalahan tersebut serta performansi dari klasifikasi algoritma tersebut. Penelitian ini menggunakan perbandingan dataset untuk data latih dan data uji sebesar 80:20, 70:30, 60:40 untuk membandingkan performansi nilai *accuracy* yang dihasilkan. Perbandingan dataset dilakukan untuk mengetahui nilai *accuracy* terbaik dari ketiganya. Hasil dari penelitian diperoleh nilai *accuracy* tertinggi didapatkan klasifikasi Multinomial Naïve Bayes Classifier dengan perbandingan dataset 80:20 sebesar 85%.

Kata kunci : *naïve bayes classifier, analisis sentimen, ulasan, aplikasi, accuracy*

1. PENDAHULUAN

Di era teknologi yang semakin canggih mendorong semua aspek bidang kehidupan ikut dalam perkembangannya, salah satunya pada bidang makanan atau kuliner. Kecanggihan hal tersebut ialah membuat perusahaan makanan berlomba-lomba untuk membuat aplikasi *mobile* untuk memberi kemudahan pada masyarakat terhadap pembelian produk yang mereka jual. Salah satu bentuk implementasi nyata dari hal tersebut dengan adanya Aplikasi My F&B ID. Aplikasi My F&B ID adalah aplikasi keluaran PT. Foods Beverages yang menaungi brand kuliner seperti Chatime, Go! Go! Curry, Gindaco dan lain-lain. Aplikasi ini dapat diunduh oleh masyarakat melalui *Google Playstore* dan *AppStore*.

Pada *Google Play Store* selain dapat mengunduh aplikasi, masyarakat juga dapat memberi pendapat terhadap aplikasi tersebut. Pendapat yang berupa memberi ulasan masukkan, kritik, pujian, keluhan terhadap aplikasi, layanan, tampilan dan menu dari aplikasi yang digunakan. Hal tersebut dapat dijadikan pertimbangan dalam mengembangkan, menganalisis serta memperbaiki aplikasi [1]. Ulasan Aplikasi dapat diberi *score rating* 1 hingga 5 oleh masyarakat untuk mengetahui kepuasannya atas penggunaan aplikasi tersebut. Namun, jika hanya dilihat melalui *score rating*, tingkat kepuasan dirasa masih kurang tepat.

Analisis sentimen dapat digunakan pada opini atau pendapat dari berbagai sumber. Selain itu, analisis sentimen dapat digunakan untuk opini sosial media seperti Twitter, Instagram, Facebook dan lain-lain. Opini dapat diklasifikasikan ke dalam kelas sentimen positif dan negatif [2]. Dalam penerapannya, analisis

sentimen memiliki dua metode yaitu metode *machine learning* dan *knowledge based method*. *Machine learning* yang digunakan dalam penelitian analisis sentimen yaitu *Support Vector Machine* [3], *Fuzzy Neural Networks* [4], *Naïve Bayes Classifier* [5], *Particle Swarm Optimization* [6] dan lain-lain. Salah satu jenis *knowledge based* adalah *Lexicon Based* [7].

Penelitian terhadap analisis sentimen ulasan aplikasi *Google Playstore* pernah dilakukan oleh Muhammad Zaki Hariansyah dan Siswanto yang menganalisis ulasan aplikasi *Livin By Mandiri* menggunakan Multinomial Naïve Bayes [8]. Hasil akhir pada penelitian tersebut berhasil mendapatkan hasil performansi yang sangat baik, dengan memperoleh nilai *accuracy* sebesar 93%, *precision* 90%, *recall* dan *f1-score* 91% serta diperoleh juga bahwa mayoritas masyarakat mengeluarkan ulasan positif terhadap aplikasi tersebut.

Penelitian lain yang membahas penggunaan Naïve Bayes sebagai metode yang digunakan dalam melakukan analisis sentimen terhadap aplikasi Tokopedia pada tahun 2019. Hasil yang didapatkan dari penelitian ini adalah nilai *accuracy* sebesar 97,13%, nilai *precision* 1, *recall* sebesar 95,49% [9].

Penggunaan Naïve Bayes dalam penelitian sangat membantu proses klasifikasi terhadap ulasan aplikasi. Hal ini mendorong peneliti untuk melakukan penelitian dalam menganalisis sentimen ulasan masyarakat mengenai aplikasi My F&B ID. Ulasan aplikasi akan dikategorikan ke dalam kelas sentimen positif dan negatif. Data yang digunakan dalam penelitian ini berasal dari data ulasan aplikasi My F&B ID di *Google Play Store* yang berjumlah 3000 ulasan

berbahasa Indonesia. Dalam analisis sentimen yang dilakukan bertujuan untuk mengetahui kecondongan masyarakat terhadap aplikasi tersebut dan juga mengetahui hasil performansi nilai *accuracy* dari model klasifikasi yang dihasilkan dengan jumlah perbandingan dataset yang berbeda.

2. TINJAUAN PUSTAKA

Pada bagian ini, peneliti akan menjabarkan dan menjelaskan tentang teori atau metode yang digunakan dalam penelitian.

2.1. Analisis Sentimen

Analisis sentimen merupakan bentuk pemrosesan bahasa lama yang berfungsi dalam mencari sentimen masyarakat menurut produk atau suatu topik [10]. Analisis sentimen adalah salah satu bidang yang berada diantara berbagai bidang seperti *Data Mining*, *Natural Language Processing* dan *Machine Learning* yang berfokus untuk ekstraksi sentimen sebuah teks [11]. Analisis sentimen bagaikan tambang pendapat yang terlibat dalam membangun sistem sebagai tempat mengumpulkan pendapat atau ulasan tentang topik atau produk yang dibangun [12].

2.2. My F&B ID

My F&B ID merupakan aplikasi perkembangan dari aplikasi sebelumnya yaitu "Chatime Indoneisa". Aplikasi ini dikeluarkan oleh PT. Food and Beverages Indonesia yang bergerak di bidang makanan dan minuman. Brand yang dijual atau yang dapat dibeli masyarakat pada aplikasi ini antara lain adalah Chatime, Cupbop, Gindaco dan lain-lain. Dengan adanya aplikasi My F&B ID, masyarakat dapat memesan secara pesan antar atau dengan memesan terlebih dahulu lalu mengambil sendiri tanpa antri dan tanpa *driver*. Pada aplikasi My F&B ID, pelanggan bukan hanya dapat memesan produk Chatime saja tetapi juga dapat memberi penilaian terhadap aplikasi.

2.3. Preprocessing Text

Preprocessing adalah tahapan untuk mempersiapkan dataset mentah untuk siap diproses [13]. Tahapan pada preprocessing meliputi sebagai berikut.

a. Case Folding

Pada tahap *case folding*, kalimat ulasan dari dataset yang memiliki huruf kapital akan diubah menjadi huruf kecil semua agar terdeteksi dengan setara atau arti yang sama.

b. Cleaning

Pada proses *cleaning* dilakukan penghilangan atau penghapusan tanda baca yang tidak diperlukan. Selain itu, pada tahap ini karakter yang tidak relevan.

c. Tokenization

Tokenization dilakukan dengan mengubah kalimat menjadi potongan-potongan kata untuk diproses lebih lanjut.

d. Normalisasi

Normalisasi digunakan untuk mengubah kata yang tidak baku menjadi kata baku.

e. Stopwords Removal

Menghilangkan kata-kata yang tidak memiliki makna seperti kata hubung dan lain-lain.

f. Stemming

Tahapan *stemming* akan mengubah potongan kata yang berimbuhan menjadi kata dasar.

2.4. Term Frequency - Inverse Document Frequency (TF-IDF)

Pembobotan yang dilakukan dengan cara perhitungan scoring dengan kamus yang telah berisi score yang telah dibuat dan dicocokkan dengan kata pada dokumen [14]. TF-IDF untuk pelabelan data yang didapatkan menurut setiap term/kata yang ada dalam kalimat atau dokumen. Bobot kata semakin besar jika sering muncul pada suatu kalimat [15].

Term frequency (TF) menghitung jumlah kemunculan kata pada dataset. Perhitungan tersebut menggunakan persamaan di bawah :

$$TF = \frac{tf}{\max(f)} \quad (1)$$

Keterangan :

tf = jumlah kata t yang muncul pada dokumen (*term frequency*).

$\max(f)$ = panjang kata pada sebuah dokumen.

Berdasarkan persamaan *Term Frequency (TF)* di atas, cara menghitung jumlah kemunculan kata t di dataset adalah membagi jumlah kata t yang muncul pada dokumen (tf) dengan panjang kata pada dokumen tersebut ($\max(f)$). *Inverse Document Frequency (IDF)* merupakan jumlah dokumen yang berisi term/kata yang akan dicari dalam dokumen dataset menggunakan persamaan di bawah :

$$idf = \log \frac{D}{df} \quad (2)$$

Keterangan :

$\log$ = logaritma natural.

D = jumlah semua dokumen.

df = jumlah *term*/kata dari dokumen.

Berdasarkan persamaan *Inverse Document Frequency (IDF)* di atas, cara menghitung jumlah dokumen yang mengandung suatu kata yaitu logaritma dari rasio total dokumen (D) terhadap jumlah kata dari dokumen (df).

TF-IDF dapat diuraikan seperti persamaan di bawah ini:

$$w = TF \times IDF \quad (3)$$

$$w = ft, d \times \log \frac{D}{df(t)} \quad (4)$$

Keterangan :

W = bobot

TF = ft, d = jumlah kata t yang muncul pada dokumen d (*term frequency*).

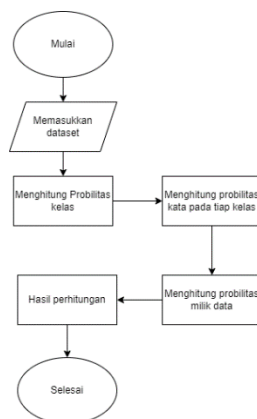
IDF = nilai *inverse document frequency*

D = jumlah kalimat atau dokumen
 $df(t)$ = jumlah dokumen yang di dalamnya terdapat kata t

Berdasarkan perhitungan kata (*TF-IDF*) di atas, cara menghitung pembobotan suatu kata pada dokumen yaitu mengalikan jumlah suatu kata pada suatu dokumen (*TF*) dengan jumlah dokumen yang mengandung kata tersebut (*IDF*).

2.5. Naïve Bayes Classifier

Naïve Bayes Classifier menjadi salah satu algoritma yang digunakan dalam analisis sentimen. Konsep dasar metode klasifikasi ini menggunakan *Teorema Bayes* yang pertama kali dinyatakan oleh Thomas Bayes [16]. Klasifikasi ini menggunakan perhitungan probabilitas terhadap kelas, fitur kemunculan fitur pada kelas serta probabilitas berdasarkan kategori pada data latih [17].



Gambar 1. Alur metode penelitian

Berdasarkan Gambar 1. Alur perhitungan algoritma Naïve Bayes yaitu dengan menghitung peluang suatu kategori pada dataset lalu menghitung peluang kata pada kategori dan selanjutnya mengklasifikasi data uji mengalikasi peluang kategori dan peluang kata pada kategori untuk didapatkan jenis kategori dari data uji. Untuk menghitung kemunculan dari peluang suatu kategori pada seluruh dokumen dapat dengan persamaan berikut [18].

$$P(c_j) = \frac{N_{c_j}}{N} \quad (5)$$

Keterangan:

$P(c_j)$ = Peluang dari kategori j

N_{c_j} = Dokumen kategori j

N = Jumlah data latih yang digunakan

Berdasarkan persamaan di atas, cara menghitung peluang kategori j ($P(c_j)$) yaitu jumlah dokumen kategori j (N_{c_j}) terhadap jumlah data latih (N).

Selanjutnya, menghitung peluang kata dalam suatu kategori. Dalam perhitungan peluang ini dibantu menggunakan konstanta α . Hal ini bertujuan untuk

agar hasil dari peluang mustahil 0, ini disebut dengan *laplace smooting*.

$$P(w_i|c_j) = \frac{\text{count}(w_i|c_j) + \alpha}{\sum w \text{count}(w|c_j) + |V|} \quad (6)$$

Keterangan :

$P(w_i|c_j)$ = Peluang sebuah kata i masuk ke kategori

j

$\text{count}(w_i|c_j)$ = total sebuah kata i dalam kategori j

$\sum w \text{count}(w|c_j)$ = total semua kata dalam kategori

j

$|V|$ = total kata unik dalam seluruh kategori

Berdasarkan persamaan di atas, cara menghitung peluang kata i kategori j ($P(w_i|c_j)$) yaitu total kata i dalam kelas kategori j ($\text{count}(w_i|c_j)$) dan $\alpha=1$ terhadap total kata dalam kategori j ($\sum w \text{count}(w|c_j)$) dan total kata unik dalam seluruh kategori ($|V|$).

Maka, proses klasifikasi pada data uji dapat dilakukan menggunakan persamaan di bawah ini.

$$P(d|c_j) = P(c_j) \times \prod_{1 \leq i \leq \text{count}(v,d)} P(w|c_j) \quad (7)$$

Keterangan :

$\text{count}(v,d)$ = total kata unik dalam dokumen

Berdasarkan persamaan di atas, cara klasifikasi terhadap data uji dalam kategori j adalah peluang kategori j dan peluang kata pada kategori j .

2.6. Confusion Matrix

Confusion matrix merupakan tabel klasifikasi kelas-kelas yang telah melalui proses algoritma yang digunakan. Bagian baris mempresentasikan jumlah data kelas actual sedangkan kolom mempresentasikan jumlah data prediksi.

1. *True Positive (TP)* = jumlah data yang diprediksi positif dan data aktualnya pun positif.
2. *False Positive (FP)* = jumlah data yang diprediksi positif namun data aktualnya negatif.
3. *False Negative (FN)* = jumlah data yang diprediksi negatif namun data aktualnya positif.
4. *True Negative (TN)* = jumlah data yang diprediksi negatif dan data aktualnya negatif.

Lalu pada confusion matrix menghitung *accuracy*, *precision*, *recall*, dan *f1-score*.

1. *Accuracy* adalah digunakan untuk mengukur tingkat keakuratan model algoritma yang digunakan dalam membuat ketetapan hasil klasifikasi.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (8)$$

2. *Precision* merupakan visualisasi untuk persentase keakuratan hasil perkiraan oleh metode yang digunakan.

$$\text{Precision} = \frac{TP}{TP+FP} \quad (9)$$

3. *Recall* adalah visualisasi dari kesesuaian metode dalam mencari ulang sebuah informasi.

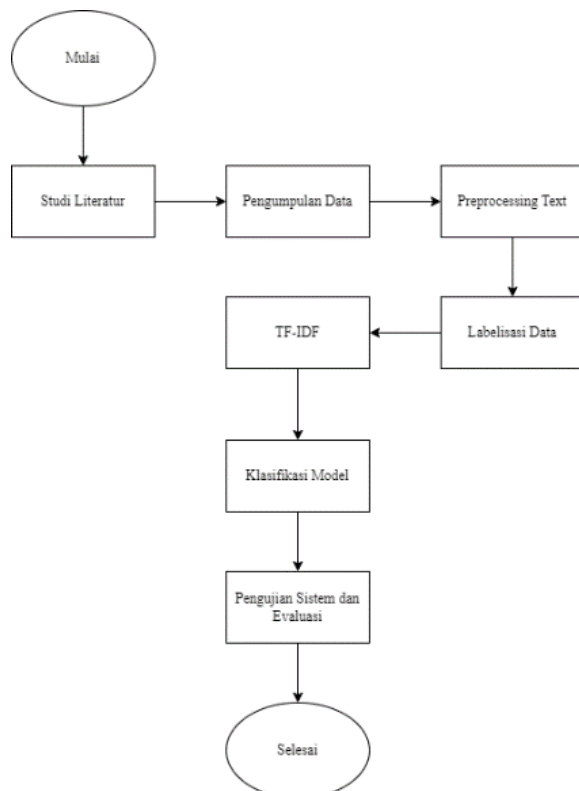
$$Recall = \frac{TP}{TP+TN} \quad (10)$$

4. *F1-score* adalah hasil perbandingan dari rata-rata nilai *precision* dan *recall* dari hasil pengujian.

$$F1 - score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (11)$$

3. METODE PENELITIAN

Pada penelitian ini menggunakan data ulasan aplikasi My FB ID dari Google Playstore. Tahapan penelitian yang akan diproses seperti bawah berikut.



Gambar 2. Alur metode penelitian

Berdasarkan gambar 2. alur metode penelitian dijabarkan sebagai berikut.

1. Mengumpulkan sumber bacaan terdahulu untuk dipelajari dan dijadikan acuan dalam penelitian yang dilakukan serta mencari teori dan metode yang dapat digunakan.
2. Mengumpulkan dataset yang berasal dari ulasan aplikasi My F&B ID dari Google Playstore sebanyak 3000 ulasan.
3. Pada tahap preprocessing, dataset yang telah dilabelisasi akan dilakukan pembersihan data agar siap menjadi dataset yang bersih untuk diklasifikasi. 6 tahapan preprocessing yang digunakan dalam penelitian ini adalah *case folding*, *cleaning*, *tokenization*, normalisasi, *stopwords removal*, *stemming*.
4. Labelisasi data dilakukan dengan kamus InSet (*Indonesian Lexicon Sentiment*) yang pada tiap kata dibuat oleh Fajri Koto dan Gemala Y.

Rahmaningtyas. Pada kamus tersebut terdapat 3.609 kata positif dan 6.609 kata negatif dan tiap kata memiliki bobot yang berbeda mulai dari -5 hingga +5[19]. Tiap ulasan akan dimasukkan ke dalam 2 kelas sentimen yaitu kelas positif dan negatif.

5. Tahap TF-IDF membobotkan tiap kata pada suatu teks yang sebanding dengan kemunculannya pada teks tersebut, namun berbanding terbalik dengan teks lain.
6. Klasifikasi model dengan pembagian dataset menjadi data latih dan data uji (80:20, 70:30, 60:40) yang digunakan pada penelitian adalah model klasifikasi Multinomial Naïve Bayes.
7. Pengujian sistem dan evaluasi dengan confusion matrix dan classification report untuk melihat hasil performansi *accuracy*.

4. HASIL DAN PEMBAHASAN

Dalam penelitian ini, dataset yang digunakan berasal dari aplikasi My F&B ID di Playstore. Pengumpulan dataset mentah berjumlah 3000 ulasan. Lalu setelah proses preprocessing dan *filtering* tambahan berupa penghapusan dokumen kosong pada dataset, data ulasan yang akan diproses untuk pelabelan terdapat 2912 ulasan. Dalam proses labelisasi menggunakan InSet tiap ulasan akan dimasukkan ke dalam kelas sentimen dan diperoleh hasil ulasan kelas positif sebanyak 1385 data dan ulasan kelas negatif sebanyak 1527 ulasan.

Hasil pengujian dari model klasifikasi Multinomial Naïve Bayes Classifier dengan 3 jenis jumlah perbandingan dataset yang telah dijelaskan sebelumnya pada bagian metodologi penelitian.

1. Skema pengujian perbandingan dataset 80:20 menghasilkan nilai *accuracy* sebesar 85%, nilai *precision* sebesar 86%, nilai *recall* sebesar 84% dan nilai *F1-score* sebesar 84%.
2. Skema pengujian perbandingan dataset 70:30 menghasilkan nilai *accuracy* sebesar 84%, nilai *precision* sebesar 86%, nilai *recall* sebesar 83% dan nilai *F1-score* sebesar 83%.
3. Skema pengujian perbandingan dataset 60:40 menghasilkan nilai *accuracy* sebesar 83%, nilai *precision* sebesar 85%, nilai *recall* sebesar 83% dan nilai *F1-score* sebesar 83%.

Perbandingan nilai *accuracy* tiap pengujian dapat dilihat pada Tabel 2.

Tabel 2. Perbandingan Nilai Accuracy Pada Klasifikasi

Perbandingan Dataset	Klasifikasi Multinomial NBC tanpa Seleksi Fitur
80:20	85%
70:30	84%
60:40	83%

Berdasarkan tabel 2. hasil performansi *accuracy* tertinggi diperoleh klasifikasi Multinomial Naïve

Bayes perbandingan dataset sebesar 80:20 senilai 85%. Sedangkan nilai performansi *accuracy* terendah diperoleh klasifikasi Multinomial Naïve Bayes perbandingan dataset sebesar 60:40 senilai 83%.

5. KESIMPULAN DAN SARAN

Hasil klasifikasi analisis sentimen ulasan aplikasi My F&B ID dengan 3 jenis pembagian dataset untuk pengujian menghasilkan performansi yang cukup baik. Dan dapat dilihat pada pembagian data tes dan data uji yang dilakukan, semakin kecil nilai data uji dan semakin besar nilai data latih maka semakin baik nilai *accuracy* yang didapatkan. Nilai *accuracy* yang diperoleh pada klasifikasi penelitian berdasarkan perbandingan data latih dan data uji, 80:20 mendapatkan 85%, 70:30 mendapatkan 84% dan 60:40 mendapatkan 83%. Untuk klasifikasi yang mendapatkan nilai *accuracy* tertinggi diperoleh klasifikasi Multinomial Naïve Bayes dengan perbandingan dataset 80:20 (2329 data latih dan 583 data uji) sebesar 85%. Lalu untuk penelitian selanjutnya diharapkan untuk dapat menggunakan jenis algoritma lain untuk analisis sentimen ulasan aplikasi atau mencoba menggunakan seleksi fitur dalam penelitian untuk menunjang model klasifikasi yang digunakan dalam penelitian.

DAFTAR PUSTAKA

- [1] Putri Nirwandani, E., & Cahya Wihandika, R. (2021). *Analisis Sentimen Pada Ulasan Pengguna Aplikasi Mandiri Online Menggunakan Metode Modified Term Frequency Scheme Dan Naïve Bayes* (Vol. 5, Issue 3). <http://j-ptiik.ub.ac.id>
- [2] Ikegami, A., Dewa, I., Bayu, M., & Darmawan, A. (2022). *Analisis Sentimen dan Pemodelan Topik Ulasan Aplikasi Noice Menggunakan XGBoost dan LDA*. In JNATIA (Vol. 1, Issue 1).
- [3] Diandra Audiansyah, D., Eka Ratnawati, D., & Trias Hanggara, B. (2022). *Analisis Sentimen Aplikasi MyXL menggunakan Metode Support Vector Machine berdasarkan Ulasan Pengguna di Google Play Store* (Vol. 6, Issue 8). <http://j-ptiik.ub.ac.id>
- [4] Anggraeny, F. T., Purbasari, I. Y., & Hatta, R. H. (2011). *PROBABILISTIC FUZZY NEURAL NETWORK UNTUK DETEKSI DINI PENYAKIT JANTUNG KORONER*. In Jurnal Informatika Mulawarman (Vol. 6, No. 2).
- [5] Darwis, D., Siskawati, N., & Abidin, Z. (n.d.). *Penerapan Algoritma Naïve Bayes untuk Analisis Sentimen Review Data Twitter BMKG Nasional*. 15(1).
- [6] Rizki, A. M., & Nurlaili, A. L. (2021). *Algoritme Particle Swarm Optimization (PSO) untuk Optimasi Perencanaan Produksi Agregat Multi-Site pada Industri Tekstil Rumahan*. *Journal of Computer, Electronic, and Telecommunication*, 1(2). <https://doi.org/10.52435/complete.v1i2.73>
- [7] Kaswidjanti, W., Himawan, H., & Silitonga, P. D. (n.d.). *The Accuracy Comparison of Social Media Sentiment Analysis Using Lexicon Based and Support Vector Machine on Souvenir Recommendations*. <http://bsd.pendukasi.id>.
- [8] Zaki Hariansyah, M. (2022). *Implementasi Metode Multinomial Naïve Bayes pada Analisis Sentimen Terhadap Layanan Aplikasi Livin by Mandiri Implementation of Naïve Bayes Multinomial Method on Sentiment Analysis of Livin by Mandiri Application Services*. In Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (SENAFTI) Jakarta-Indonesia. <https://senafiti.budiluhur.ac.id/index.php>
- [9] Apriani, R., Gustian, D., Program, S., Sistem, I., Putra, U. N., Indonesia, S., Raya, J., Kaaler, C., 21, N., & Sukabumi, K (2019). *ANALISIS SENTIMEN DENGAN NAÏVE BAYES TERHADAP KOMENTAR APLIKASI TOKOPEDIA*. In Jurnal Rekayasa Teknologi Nusa Putra (Vol 6, Issue 1).
- [10] Nuris, N., Rini Yulia, E., Solecha, K., Bina, U., Infomatika, S., & Mandiri, U. N. (2021). *Implementasi Particle Swarm Optimization (PSO) Pada Analisis Sentimen Review Aplikasi Halodoc Menggunakan Algoritma Naïve Bayes*. *Jurnal Teknologi Informasi*, 7.
- [11] A. A. K. Y. J. M. N.-K. Mahmoud Al-Ayyoub, "A comprehensive survey of arabic sentiment analysis," *Information Processing and Management*, vol. 56, no. 2, pp. 320-342, 2019.
- [12] Nuris, N., Rini Yulia, E., Solecha, K., Bina, U., Infomatika, S., & Mandiri, U. N. (2021). *Implementasi Particle Swarm Optimization (PSO) Pada Analisis Sentimen Review Aplikasi Halodoc Menggunakan Algoritma Naïve Bayes*. *Jurnal Teknologi Informasi*, 7.
- [13] Petty Wahyuningtyas, N., Eka Ratnawati, D., & Yudi Setiawan, N. (2023). *Root Cause Analysis (RCA) berbasis Sentimen menggunakan Metode K-Nearest Neighbor (K-NN)* (Studi Kasus: Pengunjung Kolam Renang Brawijaya) (Vol. 7, Issue 5).
- [14] Darwis, D., Shintya Pratiwi, E., Ferico, A., & Pasaribu, O. (2020). *PENERAPAN ALGORITMA SVM UNTUK ANALISIS SENTIMEN PADA TWITTER KOMISI PEMBERANTASAN KORUPSI REPUBLIK INDONESIA*. In IJurnal Ilmiah Edutic (Vol, 7, Issu 1).
- [15] Fremmuzar, P., & Baita, A. (2023). *Uji Kernel SVM dalam Analisis Sentimen Terhadap Layanan Telkomsel di Media Sosial Twitter*. *Komputika : Jurnal Sistem Komputer*, 12(2), 57–66.
- [16] Aldrich, J. (2008). *A. Fisher on Bayes and Bayes' Theorem*. In *Bayesian Analysis* (Vol. 3, Issue 1).
- [17] Gunawan, F., Fauzi, M. A., & Adikara, P. P. (2017). *Analisis Sentimen Pada Ulasan Aplikasi*

Mobile Menggunakan Naive Bayes dan Normalisasi Kata Berbasis Levenshtein Distance (Studi Kasus Aplikasi BCA Mobile) (Vol. 1, Issue 10).

[18] *irbookonlinereading*. (n.d.).

[19] Fajri Koto, and Gemala Y. Rahmaningtyas "InSet Lexicon: Evaluation of a Word List for Indonesian Sentiment Analysis in Microblogs". IEEE in the 21st International Conference on Asian Language Processing (IALP), Singapore, December 2017.