

COMPARISON OF K-NEAREST NEIGHBORS AND NAÏVE BAYES CLASSIFIER ALGORITHMS IN SENTIMENT ANALYSIS OF USER REVIEWS FOR INTERMITTENT FASTING APPLICATIONS

Muhammad Varhan Kusuma^{*1}, Safitri Juanita^{*2}

^{1,2}Information Systems, Faculty of Information Technology, Universitas Budi Luhur, Indonesia
Email: ¹2012500233@student.budiluhur.ac.id, ²safitri.juanita@budiluhur.ac.id

(Article received: June 21, 2024; Revision: July 10, 2024; published: October 28, 2024)

Abstract

Applications that focus on health, especially obesity prevention, are scattered in the Google Play Store, one of which is the "Intermittent Fasting" application, which, according to the developer, aims to help users maintain a healthy lifestyle and regulate eating habits. With the increasing number of similar health applications, this research focuses on sentiment analysis of user reviews of "Intermittent Fasting" to find out how users respond. The purpose of this research is to find the best algorithm to analyze sentiment on user reviews on the Google Play Store against the "Intermittent Fasting" application, as well as provide recommendations for new or old users and for application developers based on the results of processing review data. The data mining methodology used in this research is CRISP-DM, using a dataset collected on user reviews on the Google Play Store for five years (2019-2024), which is annotated with three sentiment labels (positive, negative, and neutral) based on user ratings, then modeling using two algorithms K-Nearest Neighbors (KNN) and Naïve Bayes Classifier (NBC). The contribution of this research is to test, evaluate, and compare the two algorithms (KNN and NBC) using two testing models (Split and K-Fold Cross Validation) and then provide recommendations for the best algorithm. The research concludes that the NBC algorithm is superior to KNN with an accuracy value of 80%, while the KNN algorithm has an accuracy value of only 71.43%. In addition, the K-Fold Cross Validation testing model is more optimal in improving the accuracy of the algorithm's performance than the Split model.

Keywords: CRISP-DM, intermittent fasting, k-nearest neighbors, naïve bayes classifier, sentiment analysis.

PERBANDINGAN ALGORITMA K-NEAREST NEIGHBORS DAN NAÏVE BAYES CLASSIFIER PADA ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI INTERMITTENT FASTING

Abstrak

Aplikasi yang berfokus pada kesehatan khususnya pencegahan obesitas tersebar di Google Play Store, salah satunya adalah aplikasi "Intermittent Fasting" yang menurut pengembang bertujuan untuk membantu pengguna menjaga pola hidup sehat dan mengatur kebiasaan makan. Dengan semakin banyaknya aplikasi kesehatan yang sama, maka penelitian ini berfokus pada analisis sentimen ulasan pengguna aplikasi "Intermittent Fasting" yang terdapat pada Google Play Store untuk mengetahui bagaimana respon pengguna. Tujuan penelitian ini adalah menemukan algoritma terbaik untuk menganalisis sentimen pada ulasan pengguna di Google Play Store terhadap aplikasi "Intermittent Fasting", serta memberikan rekomendasi bagi pengguna baru ataupun lama dan bagi pengembang aplikasi berdasarkan hasil pemrosesan data ulasan. Metodologi penambangan data yang digunakan pada penelitian ini adalah CRISP-DM, menggunakan dataset yang dikumpulkan pada ulasan pengguna di Google Play Store selama 5 tahun (2019-2024), yang dianotasi dengan 3 label sentimen (positif, negatif, dan netral) berdasarkan rating pengguna, kemudian pemodelan menggunakan 2 algoritma *K-Nearest Neighbors* (KNN) dan *Naïve Bayes Classifier* (NBC). Kontribusi penelitian ini adalah menguji, mengevaluasi dan membandingkan kedua algoritma (KNN dan NBC) menggunakan dua model pengujian (*Split* dan *K-Fold Cross Validation*) kemudian memberikan rekomendasi algoritma terbaik. Kesimpulan penelitian adalah algoritma NBC lebih unggul dari KNN dengan nilai akurasi sebesar 80%, sedangkan algoritma KNN nilai akurasinya hanya 71.43%. Selain itu, model pengujian *K-Fold Cross Validation* lebih optimal dalam meningkatkan akurasi dari performa algoritma dibandingkan model *Split*.

Kata kunci: analisis sentimen, CRISP-DM, intermittent fasting, k-nearest neighbors, naïve bayes classifier.

1. PENDAHULUAN

Menjaga berat badan yang ideal untuk status gizi yang normal, maka dapat mencegah terjadinya obesitas [1]. Obesitas berkaitan dengan penyakit jantung coroner [2], serta dapat menyebabkan peradangan kronis sehingga berdampak pada semakin turunnya fungsi sel dan organ manusia [3].

Selain faktor kesehatan, ada juga tren kecantikan. Dimana semakin negatif *body image* maka semakin besar pula wanita akan melakukan diet [4]. Sehingga banyak wanita yang ingin memperbaiki bentuk tubuhnya dengan mengikuti pola makan, karena diyakini bahwa standar kecantikan wanita adalah kurus [5].

Terkait dengan pencegahan obesitas dan tren kecantikan, maka salah satu aplikasi yang membantu proses tersebut adalah aplikasi “Intermittent Fasting” yang merupakan salah satu aplikasi untuk menjaga pola hidup sehat dan mengatur kebiasaan makan terhadap pengguna. Aplikasi ini diluncurkan pada tahun 2019 dan diunduh lebih dari 10 juta kali dengan 817.000 ulasan pengguna di Google Play Store.

Ulasan pengguna pada aplikasi “Intermittent Fasting” mengandung sentimen negatif, positif dan netral terkait dengan fitur pada aplikasi. Ulasan tentang obesitas, diet dan olah raga pada sosial media sangatlah dinamis, jika dikaitkan dengan obesitas maka sebagian besar mengungkapkan sentimen negatif, sedangkan untuk topik obesitas, diet dan olah raga sentimen terkadang positif, negatif ataupun netral [6].

Aplikasi yang memberikan dukungan untuk pola hidup sehat selain aplikasi “Intermittent Fasting” antara lain adalah “BodyFast: Intermittent Fasting”, “Kompanion Intermittent Fasting”, dan “HealthifyMe-Calorie Counter”. Namun dengan semakin banyaknya aplikasi yang memiliki tujuan sama maka perlu dilakukan analisis sentimen terhadap ulasan pengguna pada aplikasi “Intermittent Fasting” dengan algoritma klasifikasi.

Penelitian ini mengimplementasikan dua algoritma klasifikasi untuk sentimen analisis yaitu *K-Nearest Neighbors* (KNN) dan *Naïve Bayes Classifier* (NBC), kemudian membandingkan kinerja dari dua algoritma untuk menemukan algoritma yang memiliki kinerja terbaik. Algoritma NBC merupakan salah satu teknik untuk mengevaluasi atau memvalidasi akurasi model yang dibangun berdasarkan dataset yang digunakan [7]. Sedangkan algoritma KNN ini bergantung pada label kategori yang ada dalam dokumen pelatihan yang mirip dengan dokumen uji [7].

Ada beberapa penelitian sebelumnya yang telah menerapkan algoritma NBC dan KNN untuk menganalisis sentimen pada ulasan aplikasi, antar lain analisis sentimen aplikasi Bibit dan Bareksa dengan menggunakan algoritma KNN, NBC dan Decision Tree, berdasarkan hasil yang diperoleh KNN lebih unggul dibandingkan NBC maupun Decision Tree

dengan mendapatkan akurasi sebesar 85.14% pada aplikasi Bibit dan 81.70% pada aplikasi Bareksa [7], analisis sentimen pada aplikasi Shopee dengan algoritma KNN dan NBC, dimana algoritma NBC lebih unggul dari KNN dengan mendapatkan nilai akurasi tertinggi sebesar 80% [8]. Kemudian analisis sentimen ulasan pada aplikasi *E-Government* dengan algoritma NBC mendapatkan akurasi tertinggi sebesar 89% [9].

Selain itu terdapat penelitian analisis sentimen ulasan pengguna pada aplikasi transportasi *online* Maxim dengan algoritma KNN dan NBC, hasil penelitian menemukan bahwa algoritma NBC lebih unggul dari algoritma KNN dengan nilai akurasi NBC sebesar 81.03% pada sentimen terkait aplikasi dan 94% untuk sentimen terkait layanan [10]. Penelitian lainnya analisis sentimen pada ulasan aplikasi PLN Mobile menggunakan algoritma NBC dan KNN, hasil penelitian membuktikan bahwa dengan pengujian K-fold cross validation, algoritma KNN menghasilkan akurasi tertinggi dibandingkan algoritma NBC dengan nilai akurasi sebesar 85.97% [11].

Selain itu penelitian tentang klasifikasi sentimen tentang depresi pada Youtube menggunakan algoritma NBC dan mendapatkan nilai akurasi sebesar 84.11% [12], sentimen ulasan aplikasi Shopee pada Google Play Store dengan algoritma NBC dan menunjukan nilai akurasi mencapai 81% [13], analisis sentimen ulasan aplikasi Mypertamina dengan algoritma KNN dan NBC menghasilkan masing-masing nilai akurasi sebesar 70,73% untuk Naïve Bayes dan sebesar 85,97% untuk K-Nearest Neighbor [14], serta analisis sentimen aplikasi Peduli Lindungi dengan algoritma NBC yang menghasilkan performa akurasi sebesar 83,3% [15].

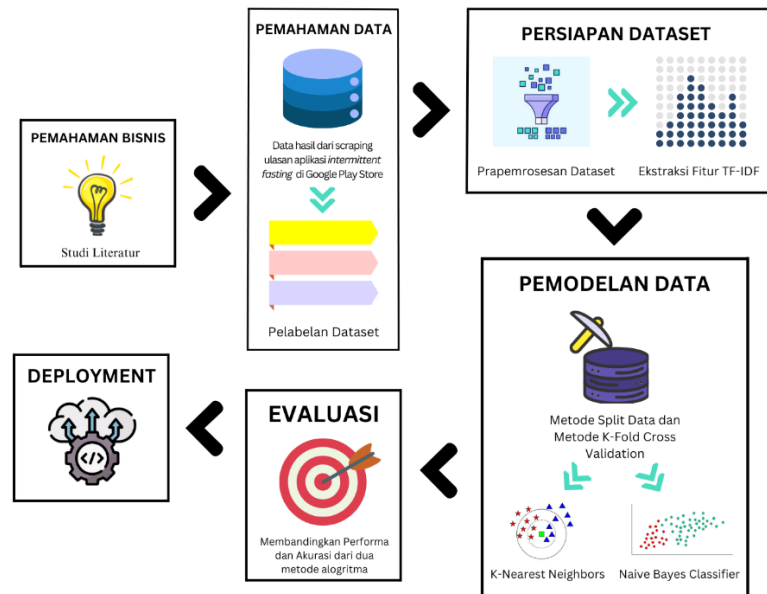
Berdasarkan hasil penelitian sebelumnya, maka dapat disimpulkan bahwa algoritma KNN dan NBC terbukti dapat diimplementasikan untuk sentimen analisis terhadap ulasan pengguna berbagai aplikasi. Sehingga kontribusi pada penelitian ini adalah sebagai berikut:

- Menggunakan data primer dengan mengumpulkan dataset yang berisi kumpulan ulasan pengguna aplikasi kesehatan bernama aplikasi “[Intermittent Fasting](#)” pada aplikasi Google Play Store selama 5 tahun serta pemberian label berdasarkan nilai rating.
- Membandingkan 2 algoritma klasifikasi (KNN dan NBC) untuk model sentimen analisis, kemudian menguji menggunakan 2 metode pengujian yang berbeda (Split dan K-Fold Cross Validation)
- Mengevaluasi, membandingkan dan menganalisis seluruh proses modeling pada kedua algoritma untuk memberikan rekomendasi algoritma terbaik.

Tujuan penelitian ini adalah menemukan algoritma terbaik untuk menganalisis sentimen pada

ulasan pengguna di Google Play Store terhadap aplikasi “Intermittent Fasting”, serta memberikan rekomendasi bagi pengguna baru ataupun lama dan

bagi pengembang aplikasi berdasarkan hasil pemrosesan data ulasan.



Gambar 1. Tahapan Proses Penelitian Perbandingan Algoritma K-Nearest Neighbors (KNN) dan Naive Bayes Classifier (NBC) pada Analisis Sentimen Ulasan Pengguna Aplikasi “Intermittent Fasting”

2. METODE PENELITIAN

Berikut merupakan tahapan dalam penelitian ini yang menggunakan metodologi penelitian CRISP-DM (*Cross-Industry Standard Process for Data mining*). CRISP-DM adalah model proses yang menjadi standar *de facto* dalam proyek *data mining* [16]. Terdapat 6 tahapan metodologi CRISP-DM yang digunakan dalam penelitian ini, pada Gambar 1 terdapat penjelasan dalam setiap tahapannya. Dalam proses eksperimen menggunakan bahasa pemrograman *Python* dan menggunakan Google Colab dan menggunakan beberapa pustaka, yaitu *Pandas*, *io*, *NLTK*, *Re*, *Numpy*, *Matplotlib*, *Seaborn*, *Sklearn* dan *WordCloud*.

2.1. Pemahaman Bisnis

Pemahaman bisnis berisikan tahapan studi literatur untuk mencari informasi mengenai referensi dari penelitian sebelumnya yang telah melakukan riset berkaitan dengan analisis sentimen menggunakan algoritma *K-Nearest Neighbors* (KNN) dan *Naive Bayes Classifier* (NBC). Algoritma KNN menentukan klasifikasi berdasarkan contoh-contoh dasar tanpa membangun representasi deklaratif eksplisit dari kategori. Algoritma ini bergantung pada label kategori yang ada dalam dokumen pelatihan yang mirip dengan dokumen uji [7]. Algoritma NBC merupakan salah satu teknik untuk mengevaluasi atau memvalidasi akurasi model yang dibangun berdasarkan dataset yang digunakan [7].

2.2. Pemahaman Data

Proses pemahaman data ini terdiri dari 2 tahapan *scraping* dan pemberian label data (*labeling*).

a. Scraping

Pada tahapan *scraping* menggunakan ekstensi dari peramban Google Chrome bernama Instant Data Scraper. Dataset merupakan ulasan pengguna di playstore Google tentang aplikasi “Intermittent Fasting” selama 5 tahun (23 November 2019 – 28 Maret 2024).

Selanjutnya pada tahapan Labeling menggunakan hasil dari rating yang dijadikan label sentimen antara positif, netral dan negatif. Pada Tabel 1 merupakan contoh dataset mentah hasil *scraping* yang digunakan pada penelitian ini.

Tabel 1. Contoh Dataset Pada Ulasan Pengguna di Playstore Google tentang Aplikasi “Intermittent Fasting”

No	Rating	Ulasan	Perangkat
1	1	<i>I bought a yearly premium , and after 3 months, I start seeing ads and my account become non premium again. This is a scam!</i>	Tablet
po2	2	<i>Ads are excessive, and things it advertises, like the body status are no where to be found.</i>	Ponsel
... 1.050	... 5	<i>... This app is great. Now I can learn how to eat less lol.</i>	... Chromebook (Laptop)

b. Pemberian Label Data

Pada tahapan pemberian label, jika rating 1 dan 2, maka diberi label sentimen negatif, rating 3 untuk

label sentimen netral, sedangkan rating 4 dan 5 untuk label sentimen positif. Setelah pelabelan akan terbentuk kolom tambahan yang disebut “Sentimen”.

2.3. Persiapan Dataset

Pada proses persiapan dataset terdapat 2 tahapan prapemrosesan dataset dan ekstraksi fitur TF-IDF. Prapemrosesan data digunakan agar gangguan pada dataset bisa berkurang dan mendapatkan hasil performa yang maksimal [17].

a. Pra-pemrosesan Dataset

Beberapa langkah untuk prapemrosesan dataset, yang pertama *Case Folding*, langkah kedua *Filtering*, yang ketiga *Tokenization*, pada tahapan ke-empat *Stop Words*, langkah kelima *Lemmatization*, dan langkah terakhir reduksi dan penggabungan kata.

b. Ekstraksi TF-IDF

Setelah pra-pemrosesan data, langkah selanjutnya mengekstraksi fitur dengan TF-IDF dan direpresentasikan dalam bentuk array. Metode *Term Frequency-Inverse Document Frequency* (TF-IDF) digunakan untuk menentukan pentingnya sebuah kata dalam suatu dokumen terhadap seluruh istilah dalam dokumen atau korpus. Metode ini menggabungkan dua ukuran: *Term Frequency* (TF) dan *Inverse Document Frequency* (IDF) [18].

2.4. Model Sentimen Analisis

Pada tahap ini menggunakan 2 model latih dan pengujian, yaitu model pengujian *Split Data* menggunakan perbandingan 80:20, 70:30, 60:40. Kemudian membandingkan dengan model *K-Fold Cross Validation*. Pada kedua model tersebut, tahap berikutnya adalah melakukan proses pemodelan dengan membandingkan performa algoritma KNN, dan NBC. Secara umum, rumus formula *Euclidean distance* dalam algoritma *K-Nearest Neighbors* sebagai berikut [19]:

$$dis(x_1, x_2) = \sqrt{\sum_{i=0}^n (x_{1i} - x_{2i})^2} \quad (1)$$

$$dis = \sqrt{\sum_{i=0}^n (x_{1i} - x_{2i})^2 + (y_{1i} - y_{2i})^2 + \dots} \quad (2)$$

Dan untuk rumus formula dalam algoritma *Naïve Bayes Classifier* sebagai berikut [19]:

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)} \quad (3)$$

Keterangan

A, B = Kejadian

P(A|B) = Probabilitas A jika B benar

P(B|A) = Probabilitas B jika A benar

P(A), P(B) = Probabilitas independen dari A dan B

2.5. Evaluasi Performa Model

Pada tahapan evaluasi membandingkan performa dan akurasi model. Terdapat dua metode pengujian, yaitu *Split Data* dan *K-Fold Cross*

Validation. Metode *Split Data* merupakan proses membagi dataset menjadi dua bagian, yaitu data latih dan data uji untuk mengevaluasi kinerja model *machine learning* [20]. Pada *Split Data* menggunakan beberapa paket dari “sklearn_metric”, yaitu *accuracy_score*, *precision*, *recall_score*, *f1_score*. Rumus formula untuk metode *Split Data* sebagai berikut [20]:

$$\{Z_i^*\}_{i=1}^n \in \operatorname{Argmin}_{Z_1, \dots, Z_n} \left\{ \frac{2}{nN} \sum_{i=1}^n \sum_{j=1}^N |z_i - z_j| \right\} \quad (4)$$

Metode *K-Fold Cross Validation* merupakan teknik umum untuk evaluasi kinerja model pada *machine learning*, teknik ini membagi data menjadi beberapa bagian seperti 3, 5, 7, 10, 15, atau 20. Di mana setiap bagian bergantian untuk digunakan sebagai data uji, sementara yang lainnya digunakan sebagai data latih. Proses ini diulang beberapa kali sehingga terdapat pada setiap bagian data menjadi data uji [21]. Pada *K-Fold Cross Validation* menggunakan dua paket dari “sklearn_model_selection” yaitu, *cross_val_score* dan *StratifiedKFold*. Rumus formula untuk metode *K-Fold Cross Validation* sebagai berikut [21]:

$$accuracy(\%) = \left(\frac{TP+TN}{(TP+FN+FP+TN)} \right) \times 100 \quad (5)$$

$$AUC = \int_0^1 \frac{TP}{(TP+FN)} d \frac{FP}{(FP+TN)} = \int_0^1 \frac{TP}{P} d \frac{FP}{N} \quad (6)$$

Keterangan:

TP = true positive

TN = true negative

FN = false negative

FP = false positive.

2.6. Deployment

Pada tahap ini, melakukan penyusunan laporan serta visualisasi hasil dari pemodelan. Terdapat pula visualisasi data menggunakan *Word Cloud* [11]. Pada tahap ini akan menampilkan kata-kata yang sering muncul dalam sentimen negatif, netral dan positif. Serta menampilkan hasil perbandingan ulasan berdasarkan perangkat yang digunakan oleh pengguna.

3. HASIL DAN PEMBAHASAN

3.1. Hasil dari Tahapan Pemahaman Data

Pada Tabel 1 merupakan contoh dataset mentah hasil scraping yang digunakan pada penelitian. Pada Tabel 2 merupakan contoh dataset yang sudah diberi label pada kolom “Sentimen”.

Tabel 2. Dataset Ulasan Setelah Proses *Labeling*

Rating	Contoh Ulasan Berdasarkan Rating	Sentimen
1	I bought a yearly premium , and after 3 months, I start seeing ads and my	Negatif

Rating	Contoh Ulasan Berdasarkan Rating	Sentimen
2	account become non premium again. This is a scam!	Negatif
3	Ads are excessive, and things it advertises, like the body status are no where to be found.	Netral
4	Nice and simple, but from time to time it stops sending me notifications	Positif
5	This app helps me stay on track for my eating periods. It also motivates me to eat less.	Positif
	This app is great. Now I can learn how to eat less lol.	Positif

Setelah dilakukan pemberian label sentimen, maka total ulasan yang digunakan pada penelitian yaitu Jumlah data ulasan yang memiliki rating 1 (169), 2 (181) kedua rating tersebut merupakan sentimen negatif (350), rating 3 yang merupakan sentimen netral (350), dan rating 4 (115), rating 5 (235) sehingga total sentimen positif (350). Masing-masing kategori sentimen memiliki total 350 ulasan.

3.2. Proses Persiapan Dataset

Pada bagian ini akan menjelaskan beberapa tahapan pada Prapemrosesan Dataset, antara lain:

1. Case Folding (mengubah teks menjadi huruf kecil). Pada Tabel 3 merupakan data yang telah melalui tahap Case Folding.

Tabel 3. Dataset Ulasan Setelah Melalui Tahap Case Folding

Contoh Ulasan Sebelum Case Folding	Contoh Ulasan Sesudah Case Folding
I bought a yearly premium , and after 3 months, I start seeing ads and my account become non premium again. This is a scam!	i bought a yearly premium , and after 3 months, i start seeing ads and my account become non premium again. this is a scam!
This app helps me stay on track for my eating periods. it also motivates me to eat less.	this app helps me stay on track for my eating periods. it also motivates me to eat less.

2. Menghapus karakter non-alfanumerik dan tanda baca (*Filtering*). Pada Tabel 4 merupakan hasil dari langkah *Filtering*.

Tabel 4. Dataset Ulasan Setelah Melalui Tahap Filtering

Contoh Ulasan Sesudah Case Folding	Contoh Ulasan Sesudah Filtering
i bought a yearly premium , and after 3 months, i start seeing ads and my account become non premium again. this is a scam!.	i bought a yearly premium and after months i start seeing ads and my account become non premium again this is a scam
this app helps me stay on track for my eating periods. it also motivates me to eat less.	this app helps me stay on track for my eating periods it also motivates me to eat less

3. Memisahkan kalimat ulasan menjadi kata-kata individual (*Tokenization*). Pada penelitian ini menggunakan paket dari "[NLTK Tokenize Punkt](#)". Pada Tabel 5 merupakan Hasil *Tokenization*.

Tabel 5. Dataset Ulasan Setelah Melalui Tahap Tokenization

Contoh Ulasan Sesudah Filtering	Contoh Ulasan Sesudah Tokenization
i bought a yearly premium and after months i start	i,bought,a,yearly,premium,and, after,months,i,start,seeing,ads,

Contoh Ulasan Sesudah Filtering	Contoh Ulasan Sesudah Tokenization
seeing ads and my account become non premium again this is a scam this app helps me stay on track for my eating periods it also motivates me to eat less	and,my,account,become,non, premium,again,this,is,a,scam this,app,helps,me,stay,on,track, for,my,eating,periods,it, also,motivates,me,to,eat,less

Pada tahapan ke-empat menghilangkan kata-kata yang sering muncul dalam kalimat ulasan tetapi tidak memiliki makna penting (*Stop Words*). Pada penelitian ini menggunakan paket dari "[NLTK StopWords](#)". Pada Gambar 2, merupakan tambahan stop words yang belum terdapat pada paket tersebut.

```
'app', 'fast', 'time', 'use', 'would', 'im', 'get', 'start', 'day', 'one',  
'way', 'need', 'open', 'hour', 'edit', 'got', 'ive', 'also', 'using',  
'many', 'version', 'water', 'much', 'times', 'fasting', 'every', 'give',  
'option', 'week', 'youre', 'see', 'end', 'days', 'eating', 'set', 'back',  
'could', 'apps', 'still', 'far', 'go', 'used', 'find', 'phone', 'timer',  
'body', 'long', 'notifications', 'going', 'started', 'make', 'add',  
'plans', 'first', 'another', 'stars', 'work', 'eat', 'hours', 'data',  
'without', 'try', 'tracking', 'custom', 'longer', 'lost', 'drink',  
'months', 'tried', 'little', 'food', 'theres', 'month', 'however',  
'intermittent', 'know', 'tracker', 'getting', 'since', 'think', 'makes',  
'manually', 'google', 'seems', 'thing', 'fasts', 'something', 'feel',  
'lot', 'year', 'allow', 'wanted', 'always', 'meal', 'instead'
```

Gambar 2. Tambahan Stop Words

Pada Tabel 6 merupakan hasil dari tahapan *Stop Words*.

Tabel 6. Dataset Ulasan Setelah Melalui Tahap Stop Words

Contoh Ulasan Sesudah Tokenization	Contoh Ulasan Sesudah Stop Words
i,bought,a,yearly,premium,and,aft er,months,i,start,seeing,ads, and,my,account,become,non, premium,again,this,is,a,scam this,app,helps,me,stay, on,track,for,my,eating, periods,it,also,motivates, me,to,eat,less	bought,yearly,premium, seeing,ads,account, become,non,premium, scam helps,stay,track,periods, motivates,less

4. Mengubah kata-kata pada kalimat ke bentuk dasarnya (*Lemmatization*). Pada penelitian ini menggunakan paket dari "[NLTK Stem WordNet](#)". Pada Tabel 7 merupakan hasil dari tahapan terakhir yaitu *Lemmatization*.

Tabel 7. Dataset Ulasan Setelah Melalui Tahap Lemmatization

Contoh Ulasan Sesudah Stop Words	Contoh Ulasan Sesudah Lemmatization
bought,yearly,premium,seeing, ads,account,become,non, premium,scam helps,stay,track,periods, motivates,less	bought,yearly,premium, seeing,ad,account, become,non,premium,scam help,stay,track,period, motivates,le

Pada Tabel 8 merupakan tahapan akhir dalam pramerosesan data, yaitu menghapus redundansi kata yang sama, lalu menggabungkan kembali kata-kata menjadi kalimat.

Tabel 8. Dataset Ulasan Setelah Redudansi dan Pengabungan Kata-kata Menjadi Kalimat

Contoh Ulasan Sesudah Lemmatization	Contoh Ulasan Sesudah Redudansi dan Penggabungan
bought,yearly,premium,seein g,ad,account,become,non, premium,scam help,stay,track,period, motivates,le	become account non scam see buy ads yearly premium periods less track stay motivate help

3.3. Tahap Pemodelan Data

Pada tahap ini akan dilakukan beberapa tahap pemodelan data, yaitu:

1. Pre-eksperimental

Pada penelitian ini, melakukan beberapa percobaan pada tahap pre-eksperimental untuk menemukan pengaturan terbaik pada kedua model (algoritma KNN dan NBC) sehingga kedua model memiliki kinerja lebih optimal dalam memodelkan data. Beberapa proses percobaannya disajikan pada Tabel 9, 10 dan 11.

Tabel 9. Hasil Nilai Akurasi dengan Proses *Stop Words*

Algoritma	K-Fold	Split		
		80:20	70:30	60:40
KNN	60.38%	63.33%	62.22%	63.57%
NBC	71.52%	69.05%	68.25%	68.33%

Tabel 10. Hasil Nilai Akurasi dengan Proses *Stop Words* + Tambahan *Stop Words*

Algoritma	K-Fold	Split		
		80:20	70:30	60:40
KNN	59.52%	60.00%	60.32%	59.76%
NBC	71.81%	68.57%	66.35%	65.95%

Tabel 11. Hasil Nilai Akurasi dengan Proses *Stop Words* + Tambahan *Stop Words* + Lemmatization

Algoritma	K-Fold	Split		
		80:20	70:30	60:40
KNN	57.62%	59.05%	60.00%	60.48%
NBC	72.10%	70.48%	66.98%	66.43%

Pada Tabel 9, pra-pemrosesan hanya menggunakan *Stop Words* yang menunjukkan hasil akurasi tertinggi pada algoritma NBC sebesar 71.52% dengan metode *K-Fold Cross Validation*, sedangkan KNN mencapai akurasi tertinggi sebesar 63.57% pada metode split 60:40.

Kemudian percobaan kedua terdapat pada Tabel 10, dimana selain menggunakan tahapan *Stop Words*, kami juga menambahkan kata-kata lain pada tahapan *Stop Words*. Penambahan kamus kata pada tahapan *Stop Words* mendapatkan kenaikan angka akurasi pada algoritma NBC menjadi 71.81% dengan metode *K-Fold Cross Validation*, sedangkan KNN mencapai akurasi tertinggi sebesar 60.32% pada metode split 70:30.

Kemudian percobaan ketiga terdapat pada Tabel 11, yaitu proses pada percobaan kedua ditambah dengan proses *lemmatization*. Pada percobaan ketiga, hasil eksperimen menunjukkan peningkatan akurasi pada algoritma NBC menjadi 72.10% dengan metode *K-Fold Cross Validation*, sedangkan KNN mencapai

akurasi tertinggi sebesar 60.48% pada metode split 60:40.

Berdasarkan 3 percobaan pada tahap pre-eksperimental penelitian ini, kami menggunakan percobaan pre-eksperimental yang terdapat pada Tabel 13 agar mendapatkan hasil yang akurat dan konsisten, serta tergoranisir dalam pemilihan kata-kata pada ulasan pengguna aplikasi "Intermittent Fasting".

2. Perbandingan Kedua Model (KNN dan NBC)

Pada tahap ini menampilkan hasil perbandingan kedua algoritma KNN dan NBC dalam melakukan analisis sentimen ulasan pada aplikasi "Intermittent Fasting" menggunakan 2 (dua) metode latih dan uji sebagai berikut:

a. Metode Split

Pengujian pada metode Split dilakukan dengan beberapa teknik pembagian data latih dan uji, dengan perbandingan 80:20, 70:30, dan 60:40. Pada Tabel 12 menampilkan hasil performa dari metrik *Accuracy*, metrik *Precision*, metrik *Recall* dan metrik *F1-Score* terhadap algoritma KNN dan NBC.

Berdasarkan Tabel 12, algoritma NBC lebih unggul dari algoritma KNN pada semua pengujian, dimana nilai akurasi tertinggi algoritma NBC adalah 70.48 pada pengujian 80:20, sedangkan algoritma KNN mendapatkan akurasi tertinggi hanya 60.48 pada pengujian 60:40.

Algoritma NBC juga unggul karena mendapatkan performa tertinggi pada pengujian split 80:20 dengan nilai precision (69.82), Recall (70.51), dan F1-Score (69.80). Sedangkan algoritma KNN mendapatkan skor performa tertinggi pada beberapa pengujian split, dimana pengujian 80:20 mendapatkan skor precision (60.32), sedangkan Split 60:40 mendapatkan skor Recall (58.76), dan F1-Score (58.11).

Tabel 12. Hasil Nilai Akurasi *Split Data* dengan Matriks Evaluasi (tulisan yang ditebalkan adalah yang memiliki performa tertinggi)

Metrik Evaluasi	SPLIT					
	80:20		70:30		60:40	
	KNN	NBC	KNN	NBC	KNN	NBC
Accuracy	59.05	70.48	60.00	66.98	60.48	66.43
Precision	60.32	69.82	60.31	66.67	60.26	67.45
Recall	58.11	70.51	58.28	66.40	58.76	66.02
F1-Score	57.24	69.80	58.05	66.01	58.11	66.12

b. Metode *K-Fold Cross Validation*

Pada metode pengujian kedua yaitu menguji performa akurasi kedua algoritma menggunakan *K-Fold Cross Validation*, dengan menggunakan data latih dan data uji yang dibagi secara acak menjadi 10 bagian. Pada Tabel 13 merupakan hasil perbandingan metode KNN dan NBC menggunakan metode latih dan uji dengan *K-Fold Cross Validation* (K) untuk analisis sentimen ulasan pengguna aplikasi "Intermittent Fasting". Berdasarkan Tabel 13 yang menampilkan hasil pengujian metrik *Accuracy*, algoritma NBC menunjukkan performa yang konsisten dan tertinggi dalam setiap pengujian nilai (K).

Selain itu, berdasarkan Gambar 6 dapat diketahui bahwa pengguna perangkat ponsel sebagian besar memberikan rating 3 pada aplikasi "Intermittent Fasting" dengan jumlah 329 ulasan. Pengguna perangkat chromebook atau laptop sebagian besar memberikan rating 5 dengan jumlah 75 ulasan. Sedangkan pengguna perangkat tablet sebagian besar memberikan rating 5 dengan jumlah ulasan sebanyak 29.

Sehingga berdasarkan Gambar 6 maka dapat disimpulkan bahwa pengguna laptop dan tablet lebih menyukai aplikasi "Intermittent Fasting" dibandingkan pengguna perangkat ponsel.

4. DISKUSI

Eksperimen ini sejalan dengan penelitian sebelumnya [9]-[15], yang secara konsisten menunjukkan bahwa NBC memberikan hasil akurasi tertinggi dalam analisis sentimen dibandingkan dengan algoritma lainnya.

Berdasarkan hasil eksperimen dalam penelitian ini, terbukti bahwa *Naive Bayes Classifier* (NBC) lebih unggul dari KNN dalam menganalisis sentimen aplikasi dengan tingkat akurasi 80%, sementara pada *K-Nearest Neighbors* (KNN) hanya mencapai 71,43%. Perbedaan penelitian ini dengan penelitian terdahulu terdapat dua hal yang utama. Pertama, penelitian ini melakukan pre-experimental terhadap dataset yang sudah melalui tahap pra-pemrosesan dan mendapatkan hasil performa terbaik pada algoritma NBC.

Kedua, penelitian ini menggunakan dua metode pengujian, yaitu *Split Data* dan *K-Fold Cross Validation*. Dari kedua metode tersebut, ditemukan bahwa *K-Fold Cross Validation* memberikan hasil akurasi yang lebih tinggi dibandingkan dengan metode *Split Data*. Berdasarkan bukti eksperimental dan literatur terdahulu, *Naive Bayes Classifier* merupakan pilihan algoritma untuk menganalisis sentimen ulasan dengan kemampuannya menangani data secara efektif dan memberikan akurasi yang tinggi.

5. KESIMPULAN

Penelitian ini melakukan analisis sentimen ulasan pengguna pada Google Play Store terhadap aplikasi "Intermittent Fasting" dengan membandingkan 2 Algoritma yaitu KNN dan NBC. Berdasarkan hasil yang diperoleh dengan metode K-Fold Validation, algoritma NBC terbukti lebih unggul dari algoritma KNN dengan nilai akurasi sebesar 80%, sedangkan KNN hanya memperoleh nilai akurasi tertinggi sebesar 71,43%.

Hasil penelitian menunjukkan bahwa algoritma NBC memiliki performa akurasi lebih baik dibandingkan KNN terhadap penelitian sentimen ulasan pengguna aplikasi "Intermittent Fasting". Pada penelitian berikutnya kami merencanakan untuk

melakukan analisis sentimen menggunakan algoritma klasifikasi lainnya seperti random forest, SVM.

DAFTAR PUSTAKA

- [1] H. E. Ardiani, T. A. E. Permatasari, and S. Sugiatmi, "Obesitas, Pola Diet, dan Aktivitas Fisik dalam Penanganan Diabetes Melitus pada Masa Pandemi Covid-19," *Muhammadiyah Journal of Nutrition and Food Science (MJNF)*, vol. 2, no. 1, Jul. 2021, doi: 10.24853/mjnf.2.1.1-12.
- [2] F. Aulia Rahman, T. Roekmantara, N. Romadhona Prodi Pendidikan Kedokteran, F. Kedokteran, and U. Islam Bandung, "Pengaruh Obesitas terhadap Kejadian Penyakit Jantung Koroner (PJK) pada Populasi Dewasa," *Bandung Conference Series: Medical Science*, vol. 2, no. 1, pp. 1002–1008, 2022, doi: 10.29313/bcsms.v2i1.1979.
- [3] S. Rose, E. R. Noer, M. Muniroh, and A. Kartini, "Literatur Review: Pembatasan energi untuk peningkatan umur panjang. Manajemen alternatif terhadap metabolik obesitas," *AcTion: Aceh Nutrition Journal*, vol. 8, no. 1, p. 139, Mar. 2023, doi: 10.30867/action.v8i1.602.
- [4] A. Valentina Millenia and A. Kurniawan, "Hubungan Antara Citra Tubuh Dengan Sikap Perempuan Terhadap Perilaku Diet," *Berajah Journal*, vol. 2, no. 2, pp. 305–314, May 2022, doi: 10.47353/bj.v2i2.93.
- [5] D. A. Putri and R. Indryawati, "Body Dissatisfaction Dan Perilaku Diet Pada Mahasiswi," *Jurnal Psikologi*, vol. 12, no. 1, pp. 88–97, Jun. 2019, doi: 10.35760/psi.2019.v12i1.1919.
- [6] C. Setiawan, "Obesitas, Olahraga, dan Diet: Analisis Sentimen pada Twitter Berbasis Analitik Big Data," *Jurnal Olahraga Kebugaran dan Rehabilitasi*, vol. 3, no. 1, pp. 71–81, 2023.
- [7] A. Dwiki, A. Putra, and S. Juanita, "Analisis Sentimen Pada Ulasan Pengguna Aplikasi Bibit Dan Bareksa Dengan Algoritma KNN," *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 8, no. 2, pp. 636–646, Jun. 2021, [Online]. Available: <http://jurnal.mdp.ac.id>
- [8] A. Oktian Permana and Sudin Saepudin, "Perbandingan algoritma k-nearest neighbor dan naïve bayes pada aplikasi shopee," *Jurnal CoSciTech (Computer Science and Information Technology)*, vol. 4, no. 1, pp. 25–32, Apr. 2023, doi: 10.37859/coscitech.v4i1.4474.
- [9] A. I. Tanggraeni and M. N. N. Sitokdana, "Analisis Sentimen Aplikasi E-Government

- Pada Google Play Menggunakan Algoritma Naïve Bayes,” *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 9, no. 2, pp. 785–795, Jun. 2022.
- [10] D. Asfi Warraihan, I. Permana, R. Novita, and A. Marsal, “Analisis Sentimen Pengguna Transportasi Online Maxim Pada Instagram Menggunakan Naïve Bayes Classifier dan K-Nearest Neighbor,” *Jurnal Media Informatika Budidarma*, vol. 7, no. 3, pp. 1134–1143, Jul. 2023, doi: 10.30865/mib.v7i3.6336.
- [11] S. Syafrizal, M. Afdal, and R. Novita, “Analisis Sentimen Ulasan Aplikasi PLN Mobile Menggunakan Algoritma Naïve Bayes Classifier dan K-Nearest Neighbor,” *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 1, pp. 10–19, Dec. 2023, doi: 10.57152/malcom.v4i1.983.
- [12] S. Mulyani and R. Novita, “Implementation Of The Naive Bayes Classifier Algorithm For Classification Of Community Sentiment About Depression On Youtube,” *Jurnal Teknik Informatika (Jutif)*, vol. 3, no. 5, pp. 1355–1361, Oct. 2022, doi: 10.20884/1.jutif.2022.3.5.374.
- [13] K. M. Elistiana, Bagus Adhi Kusuma, P. Subarkah, and H. A. Awal Rozaq, “Improvement Of Naive Bayes Algorithm In Sentiment Analysis Of Shopee Application Reviews On Google Play Store,” *Jurnal Teknik Informatika (Jutif)*, vol. 4, no. 6, pp. 1431–1436, Dec. 2023, doi: 10.52436/1.jutif.2023.4.6.1486.
- [14] D. Tualang Ksatria, Y. Yunefri, and L. F. Lhaura Van, “Analisis Sentimen Ulasan Pengguna Aplikasi Mypertamina Pada Google Playstore menggunakan K-Nearest Neighbor dan Naïve Bayes,” *Seminar Nasional Teknologi Informasi & Ilmu Komputer (SEMASTER)*, vol. 2, no. 1, pp. 213–227, 2023.
- [15] G. K. Locarso, “Analisis Sentimen Review Aplikasi Peduli Lindungi Pada Google Play Store Menggunakan NBC,” *Jurnal Teknik Informatika Kaputama (JTIK)*, vol. 6, no. 2, 2022.
- [16] C. Schröer, F. Kruse, and J. M. Gómez, “A systematic literature review on applying CRISP-DM process model,” in *Procedia Computer Science*, Elsevier B.V., 2021, pp. 526–534. doi: 10.1016/j.procs.2021.01.199.
- [17] D. Nurmalasari, I. Hermanto, and I. Ma’ruf Nugroho, “Perbandingan Algoritma SVM, KNN dan NBC Terhadap Analisis Sentimen Aplikasi Loan Service,” *Jurnal Media Informatika Budidarma*, vol. 7, no. 3, pp. 1521–1530, Jul. 2023, doi: 10.30865/mib.v7i3.6427.
- [18] H. D. Abubakar and M. Umar, “Sentiment Classification: Review of Text Vectorization Methods: Bag of Words, Tf-Idf, Word2vec and Doc2vec,” *SLU Journal of Science and Technology*, vol. 4, no. 1 & 2, pp. 27–33, Aug. 2022, doi: 10.56471/slujst.v4i.266.
- [19] D. Sandi and E. Utami, “Analisis Sentimen Publik Terhadap Elektabilitas Ganjar Pranowo di Tahun Politik 2024 di Twitter dengan Algoritma KNN dan Naïve Bayes,” *Jurnal Media Informatika Budidarma*, vol. 7, no. 3, pp. 1097–1108, Jul. 2023, doi: 10.30865/mib.v7i3.6298.
- [20] V. R. Joseph and A. Vakayil, “Split: An Optimal Method for Data Splitting,” *Technometrics*, vol. 64, no. 2, pp. 166–176, 2022, doi: 10.1080/00401706.2021.1921037.
- [21] I. K. Nti, O. Nyarko-Boateng, and J. Aning, “Performance of Machine Learning Algorithms with Different K Values in K-fold CrossValidation,” *International Journal of Information Technology and Computer Science*, vol. 13, no. 6, pp. 61–71, Dec. 2021, doi: 10.5815/ijitcs.2021.06.05.