

BAB III

METODE PENELITIAN

3.1 Objek Penelitian

Objek penelitian dalam penelitian ini adalah ulasan dari pengguna aplikasi pinjaman online AdaKami yang diambil melalui platform *AppStore*. Aplikasi AdaKami dipilih karena memiliki tingkat popularitas yang tinggi di kalangan pengguna layanan pinjaman digital, serta menyediakan banyak ulasan yang menggambarkan pengalaman langsung pengguna terhadap layanan aplikasi. Ulasan tersebut mencakup berbagai pendapat, mulai dari apresiasi terhadap kemudahan proses pengajuan hingga kritik terhadap kualitas layanan. Data ulasan yang digunakan terdiri dari komentar dan *rating* sebagai bentuk penilaian numerik. Penelitian ini bertujuan untuk menganalisis sentimen dalam ulasan tersebut, yang akan dikelompokkan menjadi sentimen positif dan negatif.

3.2 Lokasi dan Waktu Penelitian

Penelitian ini dilakukan di rumah dengan mengumpulkan data ulasan aplikasi AdaKami dari *App Store*. Penelitian berbasis data sekunder ini tidak terikat ruang fisik tertentu, memanfaatkan komputer dengan akses internet. Analisis dilakukan di tempat yang mendukung komputasi, seperti rumah atau tempat umum dengan akses internet, sehingga memungkinkan pengolahan data secara fleksibel dan efisien.

Waktu Penelitian: Penelitian ini direncanakan berlangsung selama 4 bulan, dimulai pada September hingga Desember 2024. Tahapan-tahapan penelitian dirinci sebagai berikut:

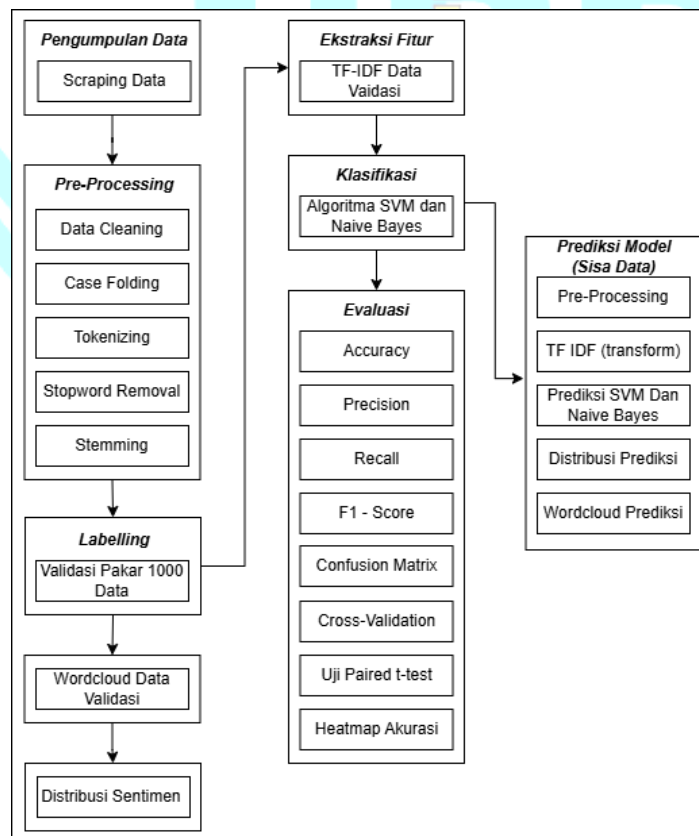
Tabel 3.1 Waktu Penelitian

No	Tahap Penelitian	Bulan 1				Bulan 2				Bulan 3				Bulan 4			
		1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4
1	Pengumpulan Data																
2	Processing Data																
3	Implementasi Model																
4	Evaluasi Model																

- 5 Analisis Hasil
Klasifikasi
- 6 Penyusunan
Laporan

3.3 Rancangan Penelitian

Penelitian ini disusun secara terstruktur melalui beberapa tahap utama, yaitu pengumpulan data, pre-processing, pembobotan, klasifikasi, dan evaluasi. Data diambil dari sumber yang relevan, kemudian diolah untuk memastikan kualitas dan kesiapannya untuk dianalisis. Selanjutnya, dilakukan pembobotan untuk menentukan tingkat signifikan dari setiap atribut. Proses klasifikasi berfungsi untuk mengelompokkan data sesuai kategori yang telah ditentukan, sementara evaluasi dilakukan untuk menilai akurasi dan efektivitas metode yang digunakan. Pendekatan ini diharapkan mampu menghasilkan data yang valid dan mendukung pencapaian tujuan penelitian.



Gambar 3.1 Rancangan penelitian

Pada Gambar 3.1 menggambarkan tahapan lengkap dalam penelitian ini, dimulai dari pengumpulan 2000 ulasan pengguna dari *AppStore* yang kemudian

diproses melalui tahapan *pre-processing*, yaitu data *cleaning*, *case folding*, *tokenizing*, *stopword removal*, dan *stemming*. Sebanyak 1000 ulasan dipilih untuk divalidasi oleh pakar dan digunakan sebagai data latih dan uji. Setelah itu, dilakukan visualisasi awal data dan ekstraksi fitur menggunakan TF-IDF. Proses klasifikasi menggunakan algoritma SVM dan *Naïve Bayes*, dengan evaluasi awal mencakup akurasi, *F1-score*, confusion matrix, dan uji paired t-test. Evaluasi lanjutan dilakukan dengan *5-fold cross-validation* untuk mengukur konsistensi model. Hasil akurasi dari cross-validation dianalisis kembali menggunakan uji *paired t-test*, serta divisualisasikan dalam bentuk *heatmap*.

Model dengan performa terbaik digunakan untuk memprediksi 971 ulasan lainnya yang belum divalidasi. Hasil prediksi ditampilkan dalam bentuk distribusi sentimen dan *wordcloud* sebagai bentuk analisis akhir.

3.4 Pengumpulan Data

Penelitian ini proses pengumpulan data menggunakan data primer yang diperoleh langsung dari sumber asli (tidak melalui perantara) yang peneliti kumpulkan secara khusus untuk memecahkan masalah penelitian. 2000 dataset yang diambil dari ulasan aplikasi AdaKami dengan kategori paling relevan, karena kategori relevan menampilkan data ulasan yang berkaitan dengan aplikasi yang tersedia pada aplikasi *App Store*, dengan cara *web scraping* menggunakan aplikasi pemrograman *python*. Hasil dari scraping data pada aplikasi AdaKami dapat dilihat pada gambar *pseudocode* di bawah ini :

Scraping
MULAI
// Instal semua alat yang diperlukan
SIAPKAN alat untuk mengambil data ulasan aplikasi
// Tentukan target aplikasi
TETAPKAN aplikasi target dengan:
- Negara: "Indonesia"
- Nama aplikasi: "adakami"
- ID aplikasi: 1462715669
// Ambil ulasan dari aplikasi
AMBIL 2000 ulasan dari aplikasi target
// Tampilkan setiap ulasan
UNTUK setiap ulasan yang diambil LAKUKAN
TAMPILKAN isi ulasan
AKHIR UNTUK
// Simpan ulasan ke dalam file tabel
KONVERSI ulasan ke format tabel
SIMPAN tabel sebagai file CSV dengan nama "adakami_reviews_id.csv"
// Unduh file hasil
UNDUH file CSV
SELESAI

Gambar 3.2 Pseudocode scraping

Hasil :

Tabel 3.2 Contoh hasil *crawling*

	Date	Review	Rating	Value
0	2020-08-31 01:32:47	Udah bunga gede, belum jatuh tempo udah di kata....	1	Negatif
1	2020-02-07 07:52:55	Jatuh tempo tgl 26 belum ada telat sama sekali...	1	Negatif
2	2022-08-06 06:31:01	Penyebaran Data Pribadi seperti Informasi Elek...	1	Negatif

1.5 Pre-processing

Tahap berikutnya merupakan *preprocessing*, pada langkah ini direncanakan jadi informasi yang bisa dianalisis. Informasi wajib diproses terlebih dulu buat membenarkan mutu informasi yang bagus saat sebelum dipakai buat analisis. Jenjang *PreProcessing* mempunyai sebagian jenjang mulai dari, *Case Folding*, *Tokenizing*, *Stopwords Removal*, *Stemming*.

3.5.1 Case Folding

Case Folding sebuah proses untuk mengubah huruf kapital menjadi huruf biasa atau standar. Proses ini mempermudah pencarian, dikarenakan tidak semua dokumen teks konsisten dengan huruf kapital. Hasil dari *case folding* dapat dilihat pada gambar *pseudocode* di bawah ini

Case Folding
MULAI #Langkah 1: Ambil data ulasan yang ada BACA data ulasan dari file atau sumber lainnya #Langkah 2: Untuk setiap ulasan dalam data UNTUK SETIAP ulasan DI data ulasan # Langkah 3: Ubah teks ulasan menjadi huruf kecil UBAH ulasan MENJADI huruf kecil #Langkah 4: Tampilkan hasil setelah case folding TAMPILKAN ulasan sebelum dan setelah perubahan SELESAI

Gambar 3.3 Pseudocode Case Folding

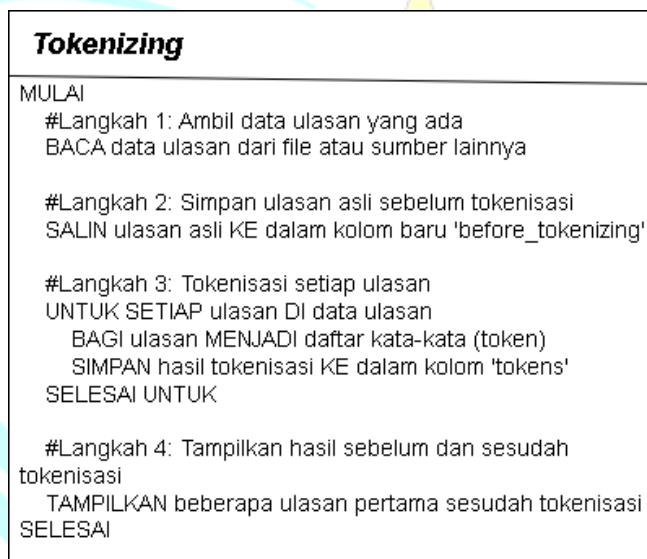
Hasil :

Tabel 3.3 Contoh Hasil *Case Folding*

Teks	Case Folding
Saya udah download udah menginstal tapimalah...	saya udah download udah menginstal tapi malah

3.5.2 Tokenizing

Tokenizing sebuah proses untuk memecahkan kalimat menjadi kata-kata yang akan menjadikan kalimat lebih bermakna.



Gambar 3.4 Pseudocode Tokenizing

Hasil :

Tabel 3.4 Contoh Hasil Tokenizing

Teks	Tokenizing
aplikasi ini keren	Aplikasi, ini, keren

3.5.3 Stopword Removal

Stopword Removal proses membuang kata-kata kurang penting atau menyimpang dari kosa kata yang tidak memiliki arti.

Stopword Removal
MULAI #Langkah 1: Ambil data ulasan yang sudah ditokenisasi BACA data ulasan yang berisi daftar token dari file atau sumber lainnya #Langkah 2: Definisikan daftar stop words TETAPKAN stop_words sebagai daftar kata-kata umum yang akan dihapus #Langkah 3: Simpan token sebelum penghapusan stop words SALIN token asli KE dalam kolom baru 'tokens_sebelum_stopwords' #Langkah 4: Hapus stop words dari setiap daftar token UNTUK SETIAP daftar token DI data ulasan HAPUS setiap kata yang ada di dalam daftar stop_words SIMPAN hasil token yang telah dibersihkan KE dalam kolom 'tokens' SELESAI UNTUK #Langkah 5: Tampilkan hasil sebelum dan sesudah penghapusan stop words TAMPILKAN beberapa ulasan pertama sebelum dan sesudah penghapusan stop words SELESAI

Gambar 3.5 Pseudocode Stopword Removal

Hasil :

Tabel 3.5 Contoh Hasil Stopword Removal

Teks	Stopword Removal
Saya sedang belajar pemrograman di kafe bersama teman-teman.	belajar pemrograman kafe teman-teman

3.5.4 Stemming

Stemming buat memperkecil jumlah indikator yang berlainan dari satu informasi alhasil tutur yang mempunyai suffix ataupun prefix hendak balik ke wujud dasarnya.

Stemming
MULAI #Langkah 1: Ambil data ulasan yang sudah dibersihkan dan ditokenisasi BACA data ulasan yang berisi daftar token dari file atau sumber lainnya #Langkah 2: Inisialisasi objek stemming TETAPKAN stemmer sebagai objek yang melakukan stemming pada kata-kata #Langkah 3: Simpan token sebelum proses stemming SALIN token asli KE dalam kolom baru 'tokens_sebelum_stemming' #Langkah 4: Lakukan stemming pada setiap token UNTUK SETIAP daftar token DI data ulasan UNTUK SETIAP kata DI daftar token Lakukan stemming pada kata SIMPAN hasil stemming KE dalam kolom 'tokens' SELESAI UNTUK #Langkah 5: Tampilkan hasil sebelum dan sesudah stemming TAMPILKAN beberapa ulasan pertama sebelum dan sesudah stemming SELESAI

Gambar 3.6 Pseudocode Stemming

Hasil :

Tabel 3.6 Contoh Hasil *Stemming*

Teks	<i>Stemming</i>
Pemain itu bermain di lapangan dan menikmati permainannya dengan penuh semangat.	main itu main di lapang dan nikmat main dengan penuh semangat

3.5 Labelling Data

Labelling merupakan tahap penting dalam analisis sentimen, di mana tanggapan masyarakat di media sosial X dikategorikan menjadi tiga kategori utama, yaitu positif, netral, dan negatif. Proses ini bertujuan untuk memberi label pada setiap entri atau komentar yang ada berdasarkan sentimen yang terkandung dalam teks tersebut. Kategori positif diberikan kepada tanggapan yang menunjukkan perasaan atau pandangan yang mendukung atau menyukai suatu topik atau kebijakan, sedangkan kategori negatif diberikan untuk tanggapan yang menunjukkan ketidaksetujuan, penolakan, atau ketidakpuasan. Kategori netral digunakan untuk tanggapan yang tidak menunjukkan sentimen yang jelas, baik mendukung maupun menentang, dan cenderung bersifat informasional atau netral.

3.6 Distribusi Sentimen

Analisis distribusi sentimen dilakukan untuk mengetahui persebaran data pada tiga kategori utama, yaitu negatif, netral, serta positif. Visualisasi disajikan dalam bentuk diagram batang guna menampilkan perbandingan jumlah antar kategori secara jelas, serta diagram lingkaran untuk memperlihatkan proporsi keseluruhan sentimen. Tujuan dari analisis ini adalah untuk mendeteksi adanya ketidakseimbangan data antar kelas yang dapat berdampak pada performa model klasifikasi.

3.7 Ekstraksi Fitur

Setelah melakukan tahapan *pre-processing* langkah selanjutnya adalah menghitung bobot dari setiap term atau kata berdasarkan frekuensi kemunculan term tersebut dalam dokumen menggunakan metode TF-IDF. TF-IDF sebuah statistik numerik yang dapat menunjukkan kata kunci dengan kata tertentu. Selain dari itu, TF-IDF juga dikenal efisien, sederhana dan akurat.

Adapun rumus pembobotan data TF-IDF sebagai berikut:

$$Wt \times d = tft, d \times idft = tft, d \times \log N / dft$$

Keterangan:

Wt, d : Bobot TF-IDF

tft, d : Jumlah frekuensi kata

$idft$: Jumlah inverse frekuensi dokumen tiap kata

dft : Jumlah frekuensi dokumen tiap kata

N : Jumlah total dokumen

3.8 Klasifikasi

Pada tahap ini, dilakukan analisis informasi berdasarkan metode yang telah ditentukan, yaitu *Naïve Bayes* dan *SVM*. Langkah ini dilakukan untuk proses klasifikasi berdasarkan nilai yang telah dianalisis sebelumnya. Hasil yang diharapkan adalah pemisahan label positif dan negatif pada ulasan aplikasi AdaKami di *AppStore* menggunakan algoritma *Naïve Bayes* dan *SVM* untuk memperoleh tingkat akurasi yang optimal (Nurkholis & Sobarnas, 2020).

Support Vector Machine (SVM) adalah algoritma yang efektif dan cepat dengan hasil yang baik untuk proses pembelajaran (Maulana, 2020). *SVM* merupakan metode pembelajaran terawasi (*supervised learning*) yang digunakan dalam klasifikasi dan analisis data. Algoritma *SVM* banyak digunakan untuk tugas klasifikasi. Dalam *SVM*, terdapat garis pemisah yang disebut *hyperplane*, yang berfungsi untuk membedakan kelas sentimen positif dan negatif. Rumus perhitungan *SVM* dapat dijelaskan sebagai berikut:

$$(w \times xi) + b = 0$$

Jika data xi termasuk dalam kelas -1, rumusnya adalah:

$$(w \times xi + b) \leq 1, yi = -1$$

Sementara itu, jika data xi termasuk dalam kelas +1, rumusnya menjadi:

$$(w \times xi + b) \geq 1, yi = +1$$

Naïve Bayes merupakan salah satu metode yang digunakan dalam data mining dengan pendekatan *supervised learning*. Metode ini didasarkan pada teorema *Bayes* dan memiliki kemampuan untuk mengklasifikasikan data dengan cara yang mirip dengan *decision tree* dan *neural network* (Rahayu, Fauzi, & Indra, 2022). Selain itu, metode *Naïve Bayes* juga memiliki waktu pemrosesan klasifikasi yang relatif cepat,

terutama dalam analisis sentimen (Maulana, 2020). Teorema *Naïve Bayes* umumnya memiliki rumus seperti berikut:

$$P(C | X) = \frac{P(X | C) \cdot P(C)}{P(X)}$$

Keterangan :

$P(C | X)$: Probabilitas kelas C untuk data X (posterior).

$P(X | C)$: Probabilitas data X muncul dari kelas C (likelihood).

$P(C)$: Probabilitas awal kelas C (prior).

$P(X)$: Probabilitas data X (evidence)

3.9 Evaluasi

Pengujian dalam penelitian ini bertujuan untuk mengevaluasi kinerja algoritma yang digunakan. Metode evaluasi yang diterapkan adalah *Confusion Matrix* (Sabily, Adikara, & Fauzi, 2020). Metode ini sangat berguna dalam menganalisis kualitas classifier. Pengujian akan menghitung nilai *accuracy*, *recall*, *precision*, dan *f1-score*, yang kemudian akan ditampilkan dalam bentuk persentase. Perhitungan *Confusion Matrix* untuk menghitung akurasi pengujian dapat dilihat pada persamaan A, persamaan B, persamaan C, persamaan D.

A. Accuracy :

Menilai sejauh mana model membuat prediksi yang tepat, dihitung dengan membagi jumlah prediksi yang benar dengan total jumlah data.

$$Accuracy = \frac{TP+TN}{TP+FN+FP+TN}$$

B. Precision :

Menghitung proporsi prediksi positif yang benar (*True Positives*) dibandingkan dengan seluruh prediksi positif yang dihasilkan oleh model.

$$Precision = \frac{TP}{TP+FP}$$

C. Recall :

Menghitung kemampuan model untuk mengidentifikasi seluruh data positif yang sesungguhnya (*True Positives*) dalam dataset.

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

D. F1-Score :

Rata-rata harmonik antara *precision* dan *recall*. Metrik ini digunakan untuk menyeimbangkan kedua nilai tersebut, terutama pada dataset yang memiliki distribusi kelas yang tidak seimbang.

$$F1 - Score = 2 \frac{Precision \times Recall}{Precision + Recall}$$

Keterangan :

TP : nilai *True Positive*

FP : nilai *False Positive*

FN : nilai *False Negative*

TN : nilai *True Negative*

Contoh Hasil Akurasi :

Tabel 3.7 Contoh Hasil Akurasi

	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Support</i>
Positif	0.85	0.90	0.87	100
Negatif	0.80	0.75	0.70	50
<i>Accuracy</i>			0.83	150

E. Confusion Matrix

Confusion matrix adalah tabel yang digunakan untuk menilai kinerja model klasifikasi dengan membandingkan prediksi model dengan data sebenarnya. Matriks ini menampilkan jumlah **prediksi** yang tepat maupun yang salah untuk setiap kategori. Sumbu vertikal matriks merepresentasikan label **aktual**, sedangkan sumbu horizontal menunjukkan label prediksi. Matriks ini berfungsi untuk memvisualisasikan pola keberhasilan dan kesalahan model dalam melakukan klasifikasi data.

Contoh Hasil *Confusion Matrix* :

Tabel 3.8 Contoh Hasil *Confusion Matrix*

Actual Class	Class		Predicted
	Positif	Negatif	
Positif	90	10	
Negatif	12	38	

F. Cross-Validation

Untuk menguji konsistensi kinerja model, digunakan teknik *k-fold cross-*

validation dengan perolehan k senilai 5. Dalam metode ini, data dipisahkan menjadi lima bagian yang digunakan secara bergiliran sebagai data uji serta data latih. Tahapan pengujian serta pelatihan dilaksanakan sebanyak lima kali, dan hasil evaluasi diperoleh dari rata-rata performa seluruh pengujian tersebut.

G. Uji *Paired t-test*

Selain metrik performa dasar, dilakukan pula pengujian *paired t-test* guna mengetahui apakah terdapat perbedaan performa yang signifikan secara statistik antara dua algoritma klasifikasi yang digunakan (SVM dan *Naïve Bayes*). Uji ini memperkuat validitas pemilihan model terbaik berdasarkan perbedaan hasil evaluasi secara matematis.

H. *Heatmap* Akurasi

Sebagai tambahan, hasil evaluasi divisualisasikan menggunakan *heatmap* untuk membandingkan performa antar model berdasarkan nilai akurasi. Visualisasi ini membantu memperjelas perbedaan kinerja secara visual dan memudahkan interpretasinya.

3.10 Prediksi Model

Setelah proses pelatihan dan evaluasi model selesai, algoritma dengan performa terbaik berdasarkan hasil pengukuran metrik dan uji statistik digunakan untuk melakukan prediksi sentimen pada 971 ulasan yang tidak melalui validasi pakar. Data ini sebelumnya telah menjalani tahapan pre-processing yang sama seperti data validasi, mencakup *stemming*, *stopword removal*, *tokenizing*, *case folding*, serta *data cleaning*.

Karena data ini tidak digunakan dalam proses pelatihan, konversi teks menjadi representasi numerik dilakukan memakai teknik *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan fungsi `transform()`, bukan `fit_transform()`, agar tidak terjadi data leakage. Proses ini memastikan bahwa model hanya menerapkan bobot kata dari hasil pelatihan sebelumnya.

Selanjutnya, model klasifikasi (SVM dan *Naïve Bayes*) diterapkan guna mengelompokkan ulasan ke dalam kategori positif, netral, atau negatif. Hasil prediksi ini tidak hanya digunakan untuk mengetahui distribusi opini pada data tanpa label, tetapi juga menjadi dasar dalam tahap visualisasi dan analisis kecenderungan opini pengguna secara umum.

Setelah klasifikasi dilakukan terhadap 971 data yang tidak divalidasi, hasil prediksi divisualisasikan dengan diagram batang yang menampilkan jumlah prediksi untuk masing-masing kategori sentimen (positif, netral, negatif) serta menggunakan *wordcloud* prediksi untuk menampilkan kata-kata dominan dari tiap kategori sentimen untuk menggambarkan tema umum dalam hasil klasifikasi.

