

BAB 3

METODE PENELITIAN

3.1 BAHAN DAN ALAT PENELITIAN

Penelitian ini menggunakan 596.000 data ulasan pengguna aplikasi WeTV yang tersedia di Google Play Store. Dikarenakan jumlah data yang sangat besar, maka diperlukan pengambilan sampel menggunakan perhitungan berdasarkan Rumus *Slovin*, sebagaimana dituliskan pada Persamaan (4). Dengan menetapkan *margin of error* sebesar 1% ($e = 0,01$) dan populasi sebesar 596.000, dengan demikian perhitungan jumlah sampel sebagai berikut:

$$n = \frac{N}{1 + N \cdot e^2} = \frac{596.000}{1 + 596.000 \cdot (0,01)^2} = \frac{596.000}{1 + 59,6} = \frac{596.000}{60,6} = 9.835$$

Berdasarkan perhitungan tersebut, diperoleh besaran *sampel* yang digunakan dalam studi ini sebanyak 9.835 ulasan. Besaran *sampel* ini dianggap memadai untuk mempertahankan akurasi hasil analisis, selaras dengan kemampuan pemrosesan data yang tersedia, serta cukup representatif dalam menggambarkan sebaran pengalaman pengguna secara proporsional.

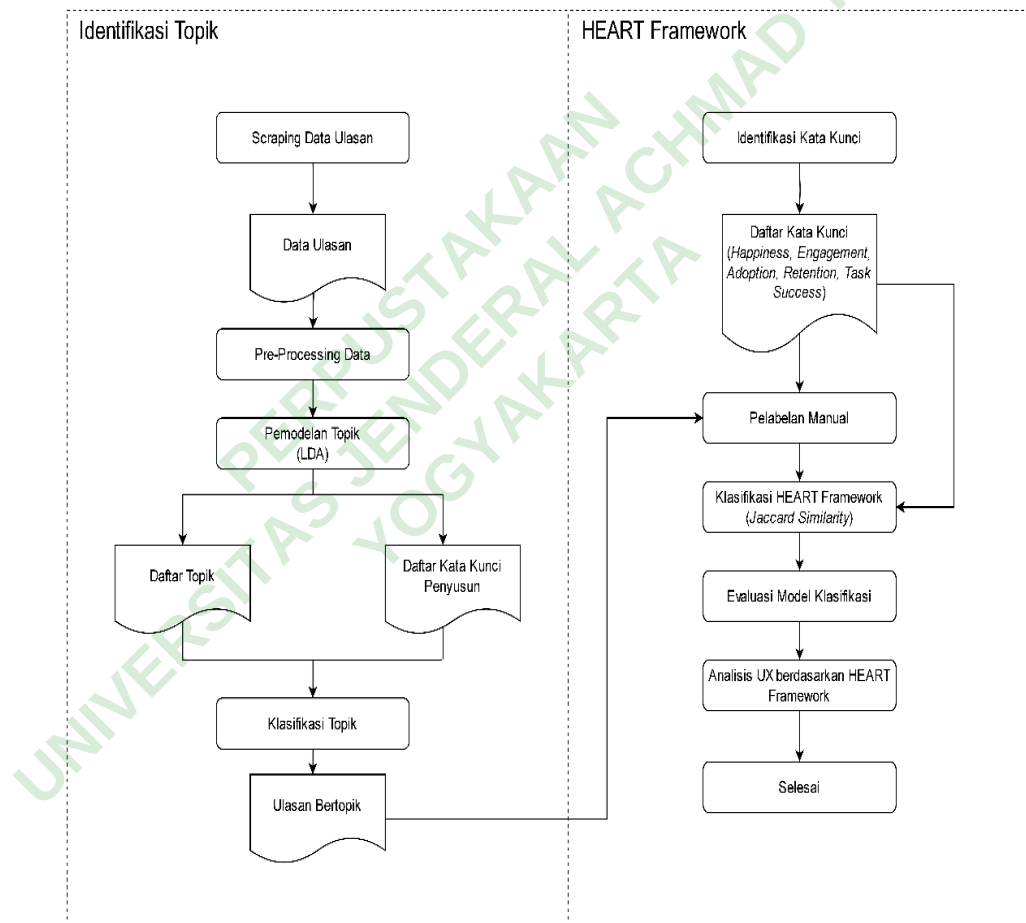
Perangkat yang digunakan mencakup *hardware*, *software*, dan sambungan internet yang cukup guna mendukung seluruh tahapan pengolahan serta analisis data. Rinciannya adalah seperti di bawah ini:

1. *Hardware* yang digunakan adalah laptop yang memiliki spesifikasi sebagai berikut:
 - a) Prosesor yang digunakan adalah Intel Celeron N4020 dengan dua inti dan frekuensi 1,10 GHz.
 - b) RAM berkapasitas 4 GB
 - c) Penyimpanan media berbasis SSD
2. *Software* yang dipakai meliputi:
 - a) *Operating system* Windows 11 versi 64-bit
 - b) *Google Colaboratory*
 - a) *Google Drive*
 - b) Pemrograman menggunakan bahasa Python

- c) *Library* yang digunakan meliputi google-play-scraper, NLTK, Sastrawi, Gensim, Numpy, Pandas, Swifter, Scikit Learn, pyLDAvis, dan Matplotlib.

3.2 JALAN PENELITIAN

Penelitian ini menerapkan metode *text mining* melalui metode pemodelan berdasarkan topik dengan LDA. Tema utama dari ulasan pengguna aplikasi WeTV diidentifikasi melalui hasil pemodelan tersebut, yang selanjutnya dianalisis secara mendalam dengan *HEART Framework* guna mengevaluasi kualitas UX, sebagaimana pada Gambar 3.1.



Gambar 3.1 Tahapan Penelitian

Langkah-langkah yang dilakukan pada penelitian ini meliputi:

1. Pengumpulan Data Ulasan melalui *Scraping*

Terdapat sekitar 596.000 ulasan dari pengguna WeTV di *platform* Google Play Store. Untuk menghitung ukuran *sampel* yang representatif dari total tersebut, digunakan Rumus *Slovin* seperti yang ditunjukkan pada Persamaan (4). Dengan margin kesalahan sebesar 1% ($e = 0,01$), diperoleh sampel sebanyak 9.835 ulasan yang kemudian disimpan dalam format .csv untuk proses analisis selanjutnya.

2. *Pre-processing*

Pada tahap ini, penulis melakukan serangkaian proses pra-pemrosesan untuk membersihkan dan mempersiapkan data ulasan agar layak dianalisis. Proses diawali dengan *cleaning*, yaitu membersihkan komponen yang tidak diperlukan seperti bilangan, simbol baca, emoji, URL, dan simbol khusus. Selanjutnya, dilakukan normalisasi untuk mengganti kata tidak baku atau singkatan ke dalam bentuk yang sesuai kaidah, misalnya “gk” menjadi “tidak” dan “udh” menjadi “sudah”. Kemudian, *case folding* diterapkan guna mengonversi seluruh karakter alfabet menjadi format huruf kecil agar menjaga konsistensi penulisan. Setelah itu, *tokenizing* memecah kalimat menjadi unit kata. Tahap berikutnya adalah penghapusan *stopword* untuk mengeliminasi kata-kata umum yang kurang bermakna seperti “yang” dan “dan”. Terakhir, proses *stemming* mengubah kata menjadi bentuk dasarnya, seperti mengubah “menonton” menjadi “tonton”. Seluruh langkah ini bertujuan agar data menjadi bersih, seragam, dan siap untuk dianalisis lebih lanjut, seperti dalam klasifikasi atau pemodelan topik.

3. Pemodelan Topik (LDA)

Langkah selanjutnya analisis ulasan dengan menerapkan algoritme LDA guna menemukan tema-tema utama yang sering muncul dalam pembahasan. LDA merupakan salah satu metode *unsupervised learning* yang berfungsi mengelompokkan teks sesuai pola kemunculan kata. Sebelum proses pemodelan dilakukan, terlebih dahulu disusun *dictionary*

dan *corpus* yang berisi daftar kata beserta frekuensi kemunculannya. Jumlah topik kemudian ditentukan berdasarkan nilai *coherence score*, yakni metrik yang merepresentasikan tingkat kekuatan hubungan antar kata dalam sebuah topik semakin besar nilainya, kualitas topik dinilai lebih baik. Tahap ini menghasilkan kumpulan topik yang berisi kumpulan kata kunci yang mewakili tema tertentu. Selanjutnya, kata tersebut dimanfaatkan sebagai dasar guna mengklasifikasikan ulasan ke dalam dimensi *HEART Framework*, dengan demikian isi ulasan dapat dihubungkan pada berbagai dimensi UX.

4. Klasifikasi Topik

Sesudah memperoleh kumpulan topik beserta kata kunci dari LDA, tiap ulasan kemudian secara otomatis dikategorikan pada topik. Pengelompokan ini didasarkan pada distribusi kata yang terdapat pada ulasan. Output dari proses ini adalah kumpulan ulasan yang sudah dilabeli dengan topik tertentu, yang nantinya akan dimanfaatkan guna mengklasifikasikan ulasan tersebut sesuai dimensi *HEART Framework*.

5. Identifikasi Kata Kunci Berdasarkan *HEART Framework*

Proses ini melibatkan penyusunan kumpulan kata kunci untuk merepresentasikan masing-masing dimensi HEART. Kumpulan kata tersebut berfungsi sebagai panduan guna mengelompokkan ulasan sesuai kesamaan maknanya. Penyusunan daftar ini merujuk pada literatur penelitian sebelumnya (Dwijayanti et al., 2022) dan disesuaikan agar relevan pada konteks aplikasi guna memastikan ketepatan tahap pengkategorian.

6. Pemberian Label secara Manual

Dalam proses ini, beberapa ulasan yang sudah melewati proses pra-pemrosesan dipilih untuk diberi label secara manual. Pemberian label dilaksanakan dengan menyesuaikan konten ulasan terhadap daftar kata kunci pada tiap dimensi HEART. Tujuannya adalah menghasilkan data berlabel yang nantinya menjadi dasar untuk mengembangkan model klasifikasi berdasarkan kesamaan kata.

7. Klasifikasi Berdasarkan *HEART Framework* menggunakan *Jaccard Similarity*

Ulasan yang belum berlabel kemudian dikategorikan secara otomatis ke salah satu dimensi *HEART* dengan memanfaatkan *Jaccard Similarity*. Metode *Jaccard Similarity* dipakai untuk menilai kesamaan antara kata pada ulasan serta kumpulan kata kunci tiap dimensi. Setiap ulasan akan mendapatkan label dari dimensi yang memiliki nilai kesamaan tertinggi. Pemilihan *Jaccard* didasarkan pada kesederhanaanya, tidak membutuhkan proses pembelajaran model yang rumit, serta efektif menilai kesamaan antar kata. Metode ini dinilai sesuai untuk data teks hasil *pre-processing* dan efisien untuk klasifikasi berdasarkan kata kunci yang digunakan untuk mengklasifikasikan ulasan ke dalam dimensi HEART.

8. Evaluasi Model Klasifikasi

Model klasifikasi berbasis kesamaan yang sudah dibuat selanjutnya diuji performanya melalui perbandingan hasil klasifikasi otomatis dan pemberian label secara manual melalui *confusion matrix*. Proses evaluasi ini menghasilkan metrik yang merepresentasikan tingkat ketepatan model dalam proses pengelompokan ulasan pengguna.

9. Analisis Pengalaman Pengguna Berdasarkan *HEART Framework*

Proses ini ulasan yang sudah dikelompokkan pada lima dimensi *HEART* dianalisis guna memahami pandangan pengguna mengenai aplikasi. Penulis mengevaluasi jumlah ulasan pada setiap dimensi dan mengaitkannya pada setiap topik yang sudah ditemukan sebelumnya. Analisis tersebut menyajikan gambaran komprehensif tentang UX berdasarkan lima dimensi HEART.

PERPUSTAKAAN
UNIVERSITAS JENDERAL ACHMAD YANI
YOGYAKARTA