

Analisis Algoritma K-Nearest Neighbor dan Naive Bayes pada Kebijakan Perkuliahan *Online* dengan Multibahasa

Analysis of K-Nearest Neighbor and Naive Bayes Algorithms on Online Lecture Policies with Multi Languages

Putri Nabila*¹, Farid Wajidi², Wawan Firgiawan³

^{1,2,3}Teknik, Teknik Informatika, Universitas Sulawesi Barat, Majene, Sulawesi Barat, Indonesia

E-mail : ptrinbila29@gmail.com¹, faridwajidi@unsulbar.ac.id²,

wawanfirgiawan@unsulbar.ac.id³

Received 25 April 2025; Revised 19 May 2025; Accepted 20 May 2025

Abstrak – Kemajuan teknologi informasi mendorong perubahan sistem pendidikan, termasuk kebijakan perkuliahan online sejak pandemi COVID-19. Penelitian ini menganalisis sentimen mahasiswa terhadap kebijakan tersebut menggunakan pendekatan multibahasa dengan algoritma *K-Nearest Neighbor* (KNN) dan *Naïve Bayes*. Adapun aspek pembeda penelitian ini adalah penggunaan data dari bahasa Indonesia dan bahasa daerah Sulawesi Barat, yaitu bahasa Mandar, yang menambah kompleksitas tahap *preprocessing*. Data dikumpulkan melalui kuesioner yang menghasilkan 1.680 opini mahasiswa, setelah tahap *preprocessing* menghasilkan 1.679 data, kemudian setelah proses pelabelan tersisa 1.280 data yang siap dianalisis. Hasil klasifikasi menunjukkan algoritma KNN terdapat 455 kategori positif dan 825 kategori negatif, serta pada algoritma *Naïve Bayes* terdapat 432 kategori positif dan 848 kategori negatif. Berdasarkan hasil klasifikasi, mayoritas opini bersifat negatif, mencerminkan bahwa mahasiswa masih kurang menerima, sehingga perlunya adaptasi dengan kebijakan tersebut. Hasil evaluasi melalui *confusion matrix* serta *K-Fold Cross Validation* menunjukkan bahwa pendekatan klasifikasi KNN mengungguli *Naïve Bayes* dengan *accuracy* 78.52%, *precision* 78.72%, *recall* 78.52%, dan *f1-score* 78.14%. Sedangkan *Naïve Bayes* mencatat *accuracy* 77.97%, *precision* 78.46%, *recall* 77.97%, dan *f1-score* 77.45%. Hal tersebut mengindikasikan bahwa KNN mempunyai kemampuan generalisasi yang lebih unggul pada konteks data teks multibahasa yang kompleks. Penelitian ini menekankan pentingnya pemilihan strategi *preprocessing* yang sesuai dalam klasifikasi sentimen multibahasa.

Kata Kunci - Analisis Sentimen, *K-Nearest Neighbor*, Multi Bahasa, *Naïve Bayes*; Perkuliahan *Online*.

Abstract – Advances in information technology have driven changes in the education system, including online lecture policies since the COVID-19 pandemic. This study analyzes student sentiment towards the policy using a multilingual approach with the *K-Nearest Neighbor* (KNN) and *Naïve Bayes* algorithms. The distinguishing aspect of this study is the use of data from Indonesian and the regional language of West Sulawesi, namely Mandar, which adds to the complexity of the preprocessing stage. Data were collected through a questionnaire that produced 1,680 student opinions, after the preprocessing stage produced 1,679 data, then after the labeling process there were 1,280 data ready to be analyzed. The classification results show that the KNN algorithm has 455 positive categories and 825 negative categories, and the *Naïve Bayes* algorithm has 432 positive categories and 848 negative categories. Based on the classification results, the majority of opinions are negative, reflecting that students are still less accepting, so it is necessary to adapt to the policy. The evaluation results through confusion matrix and *K-Fold Cross Validation* show that the KNN classification approach outperforms *Naïve Bayes* with an accuracy of 78.52%, precision of 78.72%, recall of 78.52%, and *f1-score* of 78.14%. Meanwhile, *Naïve Bayes* recorded an accuracy of 77.97%, precision of 78.46%, recall of 77.97%, and *f1-score* of 77.45%. This indicates that KNN has superior generalization

capabilities in the context of complex multilingual text data. This study emphasizes the importance of choosing the right preprocessing strategy in multilingual sentiment classification.

Keywords: *K-Nearest Neighbor, Multilingual, Naïve Bayes, Online Learning, Sentiment Analysis.*

1. PENDAHULUAN

Banyak aspek kehidupan manusia telah berubah karena kemajuan teknologi informasi, termasuk pendidikan. Sektor pendidikan mengalami transformasi besar yang semakin terasa, terutama setelah penyebaran COVID-19 di Indonesia pada tahun 2020. Kejadian tersebut mendorong pemerintah untuk segera memberlakukan berbagai kebijakan penting sebagai upaya pencegahan penyebaran virus, seperti penerapan *Work From Home* (WFH), Pembebasan Sosial Berskala Besar (PSBB), serta adaptasi kebiasaan baru (*New Normal*). Berbagai kebijakan yang diberlakukan oleh pemerintah berdampak pada dunia pendidikan, sehingga menjadi awal mula pelaksanaan perkuliahan online [1].

Sejak pertama kali diberlakukan, kebijakan perkuliahan jarak jauh (*online*) menjadi solusi utama bagi institusi pendidikan untuk tetap menjalankan proses pembelajaran di tengah situasi pandemi. Akibatnya, banyak lembaga pendidikan beralih ke metode pembelajaran jarak jauh. Kebijakan ini mengharuskan orang menggunakan internet sebagai media pembelajaran melalui *smartphone*, perangkat, atau komputer [2]. Beragam media pembelajaran daring seperti *Zoom*, *Google Meet*, *WhatsApp*, *Form Google*, *Google Classroom*, dan *E-Learning* digunakan di Indonesia untuk mengadakan perkuliahan online. Penggunaan *platform* ini semakin meningkat sejak diberlakukannya perkuliahan *online*. Survei yang dilakukan pada tahun 2020 terhadap 1.067 guru oleh Kementerian Pendidikan dan Kebudayaan (Kemendikbud) menemukan beberapa masalah utama yang menghalangi guru dan tenaga kependidikan untuk menerapkan pembelajaran *online*. Beberapa masalah tersebut termasuk manajemen pembelajaran yang buruk, pemanfaatan media digital yang buruk, serta kurangnya dorongan internal siswa untuk belajar [3].

Selain itu, terdapat hambatan lain yang sering terjadi selama pelaksanaan pembelajaran *online*, salah satunya jaringan internet tidak stabil, perangkat tidak memadai, masalah keuangan, dan masalah lainnya. Selain itu, orang tua yang kurang menguasai teknologi mungkin menghadapi tantangan dalam mendukung anak-anak mereka belajar secara *online* [4]. Hambatan yang seringkali terjadi akan menimbulkan berbagai opini yang menyebabkan kontroversi di kalangan masyarakat terkait bagaimana kesiapan Indonesia dalam mendukung pembelajaran daring (*online*).

Kebijakan terkait perkuliahan *online* memicu berbagai opini dan reaksi dari berbagai pihak, termasuk mahasiswa, dosen, guru, dan orang tua. Analisis terhadap respons masyarakat khususnya sentimen terhadap kebijakan perkuliahan *online*, menjadi semakin penting. Pemahaman yang lebih mendalam tentang pandangan dan perasaan masyarakat terhadap kebijakan ini dapat membantu institusi pendidikan dan pembuat kebijakan dalam mengembangkan sistem pendidikan di masa depan yang lebih baik. Analisis sentimen dapat digunakan untuk mengetahui seberapa efektif kuliah *online*.

Metode analisis sentimen berfungsi sebagai teknik komputasi yang digunakan agar menemukan dan mengekspresikan emosi, opini, atau penilaian dalam teks. Metode ini mengukur persentase sentimen positif, negatif, atau netral terhadap subjek tertentu, seperti individu, perusahaan, atau produk [5]. Beberapa peneliti telah mengembangkan metode analisis sentimen yang lebih kompleks dan akurat. Adapun *Naive Bayes*, *K-Nearest Neighbors*, dan *Support Vector Machine* adalah beberapa algoritma yang sering digunakan serta algoritma lainnya. Termasuk *Long-Short Term Memories* (LSTM) yang berhasil memberikan akurasi yang memuaskan [6].

Seiring dengan semakin global-nya perkuliahan *online*, tantangan utama dalam analisis sentimen terletak pada keanekaragaman bahasa yang digunakan oleh responden. Bahasa sendiri merupakan sistem komunikasi yang berfungsi untuk menyalurkan ide, emosi, serta informasi

melalui suara, tulisan, maupun simbol [7]. Data ulasan dan pendapat masyarakat tidak hanya berasal dari satu bahasa, tetapi dari beragam bahasa di berbagai wilayah. Sehingga, tujuan dari penelitian adalah untuk mengevaluasi sentimen terkait kebijakan perkuliahan *online* melalui pendekatan multibahasa yang melibatkan bahasa Indonesia, bahasa Inggris, dan bahasa Mandar sebagai bahasa daerah dari provinsi Sulawesi Barat. Data yang akan digunakan berasal dari hasil pengumpulan pendapat mahasiswa menggunakan penyebaran kuesioner

Berdasarkan masalah tersebut, penelitian ini dimaksudkan untuk menganalisis pendapat publik tentang penerapan kuliah *online* di Indonesia. Untuk melakukan analisis terkait kebijakan perkuliahan *online* diperlukan algoritma pengklasifikasian, terdapat banyak jenis algoritma pengklasifikasian. Namun dalam penelitian kali ini akan menerapkan algoritma *K-Nearest Neighbor* (KNN) dan *Naïve Bayes*. Kedua algoritma ini cukup umum digunakan dalam text mining karena prosesnya yang mudah dan cukup akurat [8].

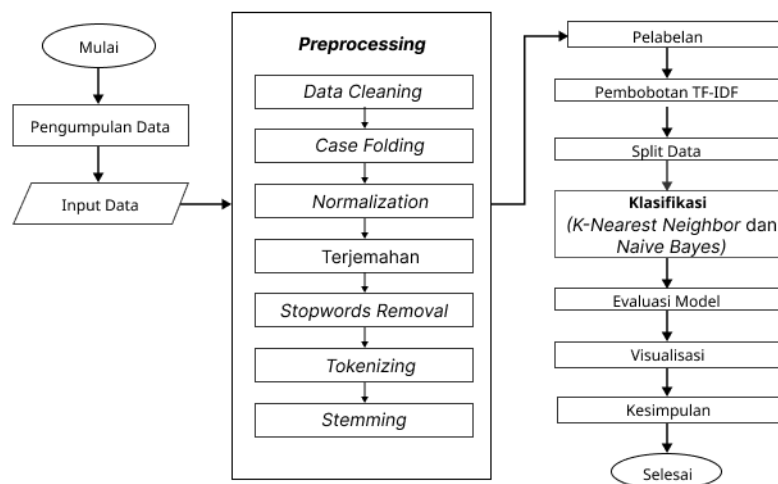
Penelitian ini didasarkan pada sejumlah studi sebelumnya yang berfokus pada analisis sentimen dengan pendekatan algoritma *Naïve Bayes* dan KNN seperti yang dilakukan oleh Andry A.P. Tangu Mara, Eko Sedyono, dan Hindriyanto Purnomo [5] dalam penelitiannya terkait pandangan mahasiswa mengenai metode pembelajaran daring yang menghasilkan algoritma KNN memperoleh tingkat 87% dan skor AUC 0,916. Selain menggunakan algoritma KNN mengenai tanggapan masyarakat mengenai pembelajaran daring, terdapat pula penelitian yang menerapkan algoritma *Naïve Bayes* yakni penelitian yang dikerjakan oleh Zulfachmi, Naufal Muhammad Kautsar, Zidan Azhari Ramadhan, Maria Brigita Mbembe, dan Yvonne [4] yang mencatat akurasi sebesar 80.67% saat menggunakan algoritma *Naïve Bayes*. Adapun penelitian lainnya yang membandingkan kedua algoritma ini yaitu penelitian oleh Silmi Nur Zaskia Wati, Herman, dan Herdianti Darwis yang menunjukkan bahwa algoritma KNN dapat diandalkan pada pengklasifikasian data, terbukti dari akurasi yang mencapai 99%, sedikit lebih tinggi dari *Naïve Bayes* dengan akurasi 98% [9]. Kemudian terdapat juga penelitian yang membandingkan 3 algoritma yakni *Naïve Bayes*, KNN dan *decision tree* pada analisis kuliah daring yang dilakukan oleh El Miana Assni Ernamia dan Asti Herliana [10]. Berdasarkan hasil analisis, algoritma *Naïve Bayes* mencatatkan tingkat akurasi sebesar 81.57%, *K-Nearest Neighbor* dengan akurasi 62.10%, dan *Decision Tree* dengan nilai akurasi 51.89%. Dari beberapa penelitian sebelumnya, penelitian ini difokuskan untuk menggunakan algoritma *Naïve Bayes* dan *K-Nearest Neighbor* yang digunakan dalam melakukan analisis sentimen mahasiswa tentang kebijakan perkuliahan *online* dengan multibahasa. Keputusan ini didasarkan pada performa algoritma *Naïve Bayes* dan KNN pada penelitian sebelumnya.

Berdasarkan latar belakang, penelitian akan menerapkan analisis sentimen terhadap kebijakan perkuliahan *online* dengan multi bahasa memanfaatkan data yang diperoleh dari penyebaran kuesioner. Tujuan dari penelitian ini adalah mengeksplorasi sentimen positif dan negatif dalam opini mahasiswa, dengan menerapkan algoritma *K-Nearest Neighbor* (KNN) dan *Naïve Bayes*.

2. METODE PENELITIAN

2.1 Tahapan Penelitian

Metodologi penelitian menguraikan urutan langkah sistematis yang diberlakukan dalam suatu penelitian untuk melaksanakan proses penelitian secara sistematis serta memastikan penelitian selaras dengan sasaran yang telah ditentukan. Rincian tahapan penelitian tersebut ditunjukkan pada gambar 1.



Gambar 1. Tahapan Penelitian

Gambar 1 menunjukkan rangkaian tahapan penelitian dalam analisis sentimen terhadap opini mahasiswa terkait kebijakan perkuliahan *online* dengan multibahasa. Setelah data dikumpulkan dari penyebaran kuesioner, dilakukan tahap *preprocessing* yang mencakup *data cleaning*, *case folding*, *normalization*, *terjemahan*, *stopword removal*, *tokenizing*, serta *stemming*. Data yang telah diproses kemudian diberi label berdasarkan kategori sentimen positif dan negatif, lalu diterapkan metode pembobotan TF-IDF diklasifikasikan menggunakan algoritma KNN dan *Naïve Bayes*. Evaluasi model dengan penerapan *k-fold cross validation* dilakukan untuk mengukur kinerja dari model yang sudah dibuat berdasarkan metrik *accuracy*, *precision*, *recall* dan *f1-score*.

2.2 Pengumpulan Data

Data utama untuk penelitian ini dikumpulkan secara langsung melalui penyebaran kuesioner terkait kebijakan perkuliahan *online*. Data yang dikumpulkan berupa opini mahasiswa mengenai kebijakan tersebut. Setelah data dikumpulkan, file tersebut akan diubah ke format *file.csv* [11].

2.3 Input Data

Setelah melakukan proses pengumpulan data dan data sudah tersedia, tahap selanjutnya adalah proses input data ke dalam sistem. Data yang akan dimasukkan mencakup pendapat mahasiswa terhadap kebijakan perkuliahan *online*. Proses ini untuk menjamin bahwa data yang diolah memenuhi persyaratan analisis sentimen, untuk mendapatkan hasil yang lebih akurat dalam memahami pandangan mahasiswa terhadap kebijakan tersebut.

2.4 Preprocessing

Tahap berikutnya dalam pengolahan teks adalah *preprocessing* yang bertujuan untuk mengubah teks mentah menjadi data numerik yang mampu ditangani oleh algoritma [12] atau mengubah data menjadi bentuk yang siap untuk dianalisis. Setelah terstruktur, data dapat diolah lebih lanjut. Pada titik ini, beberapa tindakan harus dilakukan, seperti:

- a. *Data cleaning* akan melibatkan modifikasi, penyaringan atau pengubahan. Tahap ini dilakukan untuk menghilangkan elemen-elemen yang tidak penting pada data. Hasil dari *data cleaning* yaitu semua data yang bersih dari karakter khusus seperti simbol, angka, *mention*, emotikon, tanda baca, spasi berlebih dan lainnya.

- b. *Case folding* guna mengonversi semua kata dengan huruf besar ke huruf kecil guna menghilangkan perbedaan antara huruf kapital dan huruf kecil. Ini menghasilkan setiap data yang awalnya mengandung huruf besar diubah ke huruf kecil (*lower case*).
- c. *Normalization* untuk mengubah bahasa singkatan ataupun bahasa tidak baku menjadi bahasa yang baku. Hasil dari *normalization*, yaitu semua data yang mengandung singkatan akan dinormalisasikan menjadi tidak ada lagi data yang berisi singkatan. Untuk memastikan teks konsisten dan akurat, bahasa daerah (bahasa Mandar) juga diterjemahkan ke bahasa Indonesia pada tahap ini. Koreksi kata meningkatkan kualitas data untuk analisis.
- d. Terjemahan yakni proses *translate* bahasa daerah (bahasa Mandar) ke bahasa Indonesia untuk memastikan konsistensi dan keakuratan teks. Hasil dari proses ini adalah semua data yang mengandung bahasa Mandar akan diubah menjadi bahasa Indonesia.
- e. *Stopword removal* untuk menghilangkan kata yang kurang relevan tanpa mempengaruhi pengolahan teks. Hasil dari *stopword removal*, yaitu hanya memuat kata yang akan digunakan sedangkan kata yang tidak penting dihapus.
- f. *Tokenizing* atau tokenisasi untuk mengkonversi teks menjadi bagian-bagian kecil untuk dianalisis, dikenal sebagai token. Hasil dari *tokenization*, yaitu berupa beberapa kata yang sudah dipisahkan.
- g. *Stemming* menurunkan kata ke bentuk paling mendasar atau untuk menghapus imbuhan maupun kata depan. Hasil dari *stemming*, yaitu semua kata yang berubah menjadi kata dasar.

Dari penjelasan tersebut, dapat dilihat bahwa *preprocessing* penting untuk dilakukan sebelum mengolah data agar data tersebut benar-benar sudah siap untuk diolah [13].

2.5 Pelabelan Data

Setelah melalui tahapan *preprocessing* data, selanjutnya data akan diberi label sesuai kategori. Model pembelajaran mesin berbasis transformer bernama *Bidirectional Encoder Representations from Transformers* atau biasa disebut BERT telah efektivitas tinggi dalam banyak tugas NLP, salah satunya analisis sentimen [14]. Pelabelan ini bertujuan untuk membantu dalam mengidentifikasi sentimen dalam data dengan menggunakan model BERT yakni: "nlptown/bert-base-multilingual-uncased-sentiment".

2.6 Pembobotan TF-IDF

Term Frequency – Inverse Document Frequency atau sering disebut TF-IDF digunakan untuk memberi bobot kata tertentu dinilai berdasarkan tingkat kemunculannya dalam dokumen dan keseluruhan dataset. TF-IDF memberikan bobot pada setiap kata berdasarkan kepentingannya dalam dokumen dengan mengubah kata menjadi vektor untuk klasifikasi, seperti dengan *K-Nearest Neighbor* dan *Naïve Bayes* [15]. Kata yang lebih sering ditemukan pada dokumen tertentu tetapi hanya sesekali di dataset secara keseluruhan diberi bobot lebih tinggi. Adapun pada penelitian oleh [16] untuk perhitungan pembobotan TF-IDF menggunakan persamaan berikut:

$$tf(t, d) = \frac{\text{Jumlah kemunculan kata } (t) \text{ dalam dokumen } (d)}{\text{Jumlah total kata dalam dokumen}} \quad (1)$$

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

$$w(t, d) = tf(t, d) \times idf \quad (3)$$

Rumus diatas adalah rumus untuk pembobotan menggunakan TF-IDF. Rumus tersebut termasuk perhitungan frekuensi kemunculan kata (t) muncul dalam dokumen (d), dengan t adalah kata (*term*) yang sedang dianalisis dan d adalah dokumen dimana kata tersebut muncul. Jumlah

total dokumen (N) dihitung bersama df yaitu jumlah dokumen yang mengandung kata atau t ($term$). Lalu mengalikan masing-masing hasil dari keduanya.

2.7 Split Data

Data kemudian dibagi menjadi *training set* dan *testing tes* untuk melatih serta menguji model klasifikasi. Model dilatih menggunakan *training set* agar dapat mengenali pola dalam data, sementara *testing set* digunakan untuk mengevaluasi kinerja model dalam mengklasifikasikan data baru. Pembagian ini bertujuan untuk memastikan bahwa Model bekerja dengan baik pada data latih dan mampu melakukan generalisasi pada data yang belum dikenali sebelumnya, sehingga analisis sentimen lebih akurat dan andal.

2.8 Klasifikasi

Selanjutnya tahap klasifikasi, data hasil proses pembobotan TF-IDF dipakai sebagai masukan untuk dua algoritma, yaitu *K-Nearest Neighbor* dan *Naïve Bayes*. Algoritma *K-Nearest Neighbor* yang sering disebut KNN yaitu suatu metode yang bekerja dengan mengklasifikasikan objek berdasarkan kedekatan jaraknya dengan data pelatihan yang ditempatkan pada ruang dengan jumlah dimensi yang besar, setiap dimensi menunjukkan ciri atau karakteristik dari data tersebut [17]. Sementara itu, *Naïve Bayes* merupakan algoritma klasifikasi yang didasarkan pada prinsip probabilitas sederhana, serta menghitung kemungkinan suatu kategori dengan menganalisis gabungan frekuensi dan nilai yang terdapat dalam dataset [18]. Kedua algoritma ini sering digunakan dalam *text mining* karena memiliki proses yang sederhana dan tingkat akurasi yang cukup baik [19].

2.9 Evaluasi Model

Tahap ini bertujuan untuk mengevaluasi penggunaan algoritma KNN dan *Naïve Bayes* sehingga diketahui tingkat akurasi dari keduanya. Melalui penelitian ini, digunakan metode *K-Fold Cross Validation* sebagai teknik validasi untuk menguji keakuratan model. Prinsipnya adalah melakukan perulangan sebanyak nilai K, dengan pembagian data secara acak, sehingga model dapat diuji dengan data uji dan data latih yang berbeda pada setiap iterasi. Data yang digunakan dibagi dua, yakni data pelatihan dan data pengujian, sehingga memungkinkan sistem untuk mempelajari dan menguji secara lebih adil dan menyeluruh[20].

Penilaian performa model selanjutnya dijalankan dengan bantuan *confusion matrix* berfungsi untuk menilai performa klasifikasi sentimen. *Confusion matrix* berfungsi guna menilai sejauh mana model mampu mengkategorikan data dengan menganalisis hasil antara prediksi model dan label sebenarnya. Terdapat pula metrik evaluasi yang akan dihitung seperti *accuracy*, *precision*, *recall*, dan *f1-score* [1]. Berikut adalah persamaan untuk menghitung nilai *accuracy*, *precision*, *recall*, dan *f1-score*.

$$Accuracy = \frac{(TP+TN)}{(TP+TN+FN+FP)} \quad (4)$$

$$Precision = \frac{TP}{(TP+FP)} \quad (5)$$

$$Recall = \frac{TP}{(TP+FN)} \quad (6)$$

$$F1\ Score = \frac{precision \times recall}{precision + recall} \times 2 \quad (7)$$

2.10 Visualisasi

Tujuan dari tahap visualisasi adalah untuk menyajikan ilustrasi umum mengenai kata-kata yang paling dominan dalam dataset sentimen. Pada tahap ini, data akan ditampilkan dalam

bentuk cloud kata, yang merupakan representasi visual dari teks dengan ukuran setiap kata merepresentasikan tingkat kemunculan atau beratnya dalam dataset [21]. Dalam konteks analisis sentimen, visualisasi ini meningkatkan pemahaman pola umum kata-kata yang sering diasosiasikan dengan kategori sentimen tertentu, seperti kata-kata positif dan negatif.

2.11 Kesimpulan

Tahap terakhir dalam sistem klasifikasi sentimen adalah membuat kesimpulan, di mana semua hasil dari tahap-tahap sebelumnya dirangkum untuk memberikan gambaran keseluruhan tentang kinerja sistem. Kesimpulan ini mencakup analisis dari seluruh proses, mulai dari pengumpulan data, *preprocessing*, pelabelan data, pembobotan TF-IDF, split data penerapan algoritma *K-Nearest Neighbor* (KNN) dan *Naïve Bayes*, hingga evaluasi menggunakan *confusion matrix* dan visualisasi dengan *word cloud*. Pada tahap ini, hasil evaluasi model, seperti tingkat akurasi dan pola kesalahan, dianalisis untuk menilai efektivitas model dalam mengklasifikasikan sentimen. Selain itu, kesimpulan juga dapat mencakup rekomendasi perbaikan jika model belum optimal, serta saran untuk pengembangan lebih lanjut. Dengan demikian, tahap kesimpulan menyajikan ringkasan hasil analisis dan memberikan panduan untuk langkah selanjutnya.

3. HASIL DAN PEMBAHASAN

Metode *K-Nearest Neighbor* dan *Naïve Bayes* digunakan untuk mengklasifikasikan sentimen siswa terhadap kuliah *online* serta untuk mengukur seberapa akurat hasil dari kedua algoritma dalam menilai pendapat siswa tentang kebijakan kuliah *online* multibahasa.

3.1 Pengumpulan Data

Data primer dikumpulkan secara langsung untuk penelitian ini melalui penyebaran kuesioner terkait kebijakan perkuliahan *online*. Data yang dikumpulkan berupa opini mahasiswa mengenai kebijakan tersebut, termasuk tingkat kepuasan mereka, yang dikategorikan sebagai positif atau negatif. Melalui data ini, penelitian diharapkan mampu menghasilkan temuan yang relevan dan bermanfaat sebagai acuan bagi pemerintah dan lembaga pendidikan dalam mengambil keputusan terkait kebijakan perkuliahan *online* di masa mendatang. Sampel data yang diperoleh dari penyebaran kuesioner sebanyak 1680 data.

Tabel 1. Sampel Data

No.	Timestamp	Tingkat Kepuasan Anda	Mohon Jelaskan Alasan dari Jawaban Anda Sebelumnya
1.	11/8/2024 19:41:54	Cukup Puas	puas karena lingkungan pertemanan juga sangat bagus beserta dosennya
2.	11/8/2024 19:46:06	Cukup Puas	Anna cukup puas u pilih apa' lebih malai ii tau mengerti mua' offline ii dibanding online
3.	11/8/2024 19:46:49	Tidak Puas Sama Sekali	Karena semakin mempersulit mengerti materi yang diberikan dan sosialisasi terhadap teman berkurang
4.	11/12/2024 15:37:29	Tidak Puas Sama Sekali	Bikin ngantuk Tidak fokus Waktu yang tidak teratur
5.	2/2/2025 10:12:41	Sangat Puas	Karena klau mata kuliahnya cuman 1 itu nanggung klau ke kampus

6	2/5/2025 21:43:16	Sangat Puas	karena saya bisa mengakses lebih banyak referensi dan belajar dari berbagai sumber online yang tersedia.
---	-------------------	-------------	--

Tabel 1 menunjukkan beberapa sampel dataset awal atau data mentah, yang berasal dari proses pengambilan data.

3.2 Preprocessing

Preprocessing merupakan tahap awal dari *text mining* atau *Natural Language Processing*. Data mentah akan disaring dan dibersihkan pada tahap *preprocessing* sehingga dapat digunakan untuk analisis atau pemrosesan lanjutan. Tujuannya adalah untuk mengubah data menjadi mudah dipahami dan diolah menggunakan model atau algoritma, yang akan digunakan dengan membersihkan, mengorganisir, dan mempersiapkan data. Pembersihan data dilakukan melalui berbagai proses, seperti:

3.2.1 Data Cleaning

Data cleaning adalah tahap awal pada *preprocessing* dengan tujuan untuk menghilangkan bagian yang tidak penting pada data.

Tabel 2. Tahap *Data Cleaning*

Sebelum <i>Data Cleaning</i>	Setelah <i>Data Cleaning</i>
Lebih fleksibel saja kalau kuliah online, apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan	Lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan
Pe'guruan wattu makkuliah online se iyya kurang efektif usa'ding saba' terbatasnya kuota anna' jaringan pada saat online. Diang wattu saat kuliah online juga jaringan kurang stabil moa' matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andangi mala tau berinteraksi ke sesama sola.	Peguruan wattu makkuliah online se iyya kurang efektif usading saba terbatasnya kuota anna jaringan pada saat online Diang wattu saat kuliah online juga jaringan kurang stabil moa matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andangi mala tau berinteraksi ke sesama sola
Karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus.	Karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus
NDANI EFEKTIF	NDANI EFEKTIF
Karena klau mata kuliahnya cuman 1 itu nanggung klau ke kampus	Karena klau mata kuliahnya cuman itu nanggung klau ke kampus

Tabel 2 terdapat hasil dari proses *data cleaning*. Hasil *data cleaning* mengeliminasi karakter tidak relevan seperti tanda baca, angka, emoji maupun URL. Pada opini kedua terdapat tanda koma atas (^), tanda titik (.), serta tanda koma (,), serta pada opini kelima terdapat angka yang akan dihapus, sehingga kalimat menjadi lebih bersih begitupun dengan opini lainnya.

3.2.2 Case Folding

Tahap selanjutnya adalah *case folding* yakni seluruh kata yang memiliki huruf kapital dikonversi jadi format huruf kecil.

Tabel 3. Tahap *Case Folding*

Sebelum <i>Case Folding</i>	Setelah <i>Case Folding</i>
Lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan	lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan
Peguruan wattu makkuliah online se iyya kurang efektif usading saba terbatasnya kuota anna jaringan pada saat online Diang wattu saat kuliah online juga jaringan kurang stabil moa matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andangi mala tau berinteraksi ke sesama sola	peguruan wattu makkuliah online se iyya kurang efektif usading saba terbatasnya kuota anna jaringan pada saat online diang wattu saat kuliah online juga jaringan kurang stabil moa matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andangi mala tau berinteraksi ke sesama sola
Karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus	karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus
NDANI EFEKTIF	ndani efektif
karena klau mata kuliahnya cuman itu nanggung klau ke kampus	karena klau mata kuliahnya cuman itu nanggung klau ke kampus

Pada tabel 3, terdapat opini dari mahasiswa yang merupakan hasil dari proses *data cleaning* sebelumnya. Bagian hasil dari teks yang sudah dilakukan *case folding* menunjukkan perubahan pada penggunaan huruf kecil (*lower case*), yakni pada opini ketiga terdapat kata (Karena) yang kemudian berubah menjadi (karena), pada opini keempat terdapat kata (NDANI EFEKTIF) berubah menjadi ndani efektif. Begitupun dengan opini yang lainnya.

3.2.3 Normalization

Normalization adalah tahap mengubah atau menyesuaikan data menjadi format standar atau seragam. Dapat juga berupa mengubah singkatan menjadi kata yang baku. Tujuannya adalah untuk meningkatkan konsistensi dan variasi data sehingga lebih mudah untuk membandingkannya, menganalisisnya, atau digunakan dalam berbagai konteks.

Tabel 4. Tahap *Normalization*

Sebelum <i>Normalization</i>	Setelah <i>Normalization</i>
lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan	lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan
peguruan wattu makkuliah online se iyya kurang efektif usading saba terbatasnya kuota anna jaringan pada saat online diang wattu saat kuliah online juga jaringan kurang stabil moa matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andangi mala tau berinteraksi ke sesama sola	pegguruan wattu makkuliah online se iyya kurang efektif usaqding sabaq terbatasnya kuota anna jaringan pada saat online diang wattu saat kuliah online juga jaringan kurang stabil muaq matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andiang malai tau berinteraksi ke sesama sola
karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus	karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus
ndani efektif	andiang efektif
karena klau mata kuliahnya cuman itu nanggung klau ke kampus	karena kalau mata kuliahnya cuma itu nanggung kalau ke kampus

Hasil dari proses *normalization* yang ada pada tabel 4 menunjukkan bahwa adanya perubahan dari hasil normalisasi yang menggunakan hasil dari proses *case folding* sebelumnya. Proses ini menstandarkan kata-kata agar konsisten. Misalnya, pada opini kelima kata (klau) dinormalisasikan menjadi (kalau), kata (cuman) dinormalisasi menjadi kata (cuma). Begitupun dengan kata-kata yang mengalami perubahan serupa pada opini lainnya. Kemudian hasil *normalization* dalam bahasa Mandar itu contohnya pada opini kedua, yakni kata (peguruan) yang dinormalisasikan menjadi (pegguruan), kata (usading) menjadi (usaqding), kata (saba) menjadi (sabaq), kata (moa) menjadi (muaq), serta kata (mala) menjadi (malai) sesuai dengan rujukan kamus bahasa Mandar.

3.2.4 Terjemahan

Pada tahap ini juga akan dilakukan perubahan kata dalam bahasa Mandar yang belum sesuai dengan penulisan yang baik dan benar. Tahap ini juga akan dilakukan *translate* bahasa daerah (bahasa Mandar) ke bahasa Indonesia untuk memastikan konsistensi dan keakuratan teks. Koreksi kata meningkatkan kualitas data untuk analisis. Adapun proses terjemahan ini menggunakan Kamus Bahasa Mandar-Indonesia oleh Abdul Muthalib[22].

Tabel 5. Tahap Terjemahan

Sebelum Terjemahan	Setelah Terjemahan
lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat	lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang

canggih yang dapat digunakan untuk mendukung proses perkuliahan	sangat canggih yang dapat digunakan untuk proses perkuliahan
pegguruan wattu makkuliah online se iyya kurang efektif usaqding sabaq terbatasnya kuota anna jaringan pada saat online diang wattu saat kuliah online juga jaringan kurang stabil muaq matei listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta andiang malai tau berinteraksi ke sesama sola	pelajaran waktu berkuliah online se iyya kurang efektif saya rasa sebab terbatasnya kuota dan jaringan pada saat online ada waktu saat kuliah online juga jaringan kurang stabil kalau mati listrik jadi otomatis berdampak ke jaringan dan interaksi ke dosen terbatas serta tidak dapat tau berinteraksi ke sesama teman
karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus	karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus
andiing efektif	tidak efektif
karena kalau mata kuliahnya cuma itu nanggung kalau ke kampus	karena kalau mata kuliahnya cuma itu nanggung kalau ke kampus

Hasil dari terjemahan yang ada pada tabel 5 menunjukkan semua data yang mengandung bahasa Mandar akan diubah menjadi bahasa Indonesia. Misalnya yang terdapat pada opini kedua, yakni yakni kata (pegguruan) yang diterjemahkan menjadi (pelajaran), kata (wattu) menjadi (waktu), kata (makkuliah) menjadi (berkuliah), kata (usaqding) menjadi (saya rasa), kata (sabaq) menjadi (sebab), kata (anna) menjadi (dan), kata (diang) menjadi (ada), kata (muaq) menjadi (kalau), kata (matei) menjadi (mati), kata (andiing) menjadi (tidak), kata (malai) menjadi (dapat), serta kata (sola) menjadi (teman). Begitupun opini lainnya yang mengandung kata-kata berbahasa Mandar diubah menjadi bahasa Indonesia.

3.2.5 Stopword Removal

Stopword removal digunakan untuk mengeliminasi kata-kata umum yang tidak diperlukan tanpa mempengaruhi analisis teks. *Stopword* biasanya terdiri dari kata-kata umum seperti “dan”, “atau”, “yang”, “di”, “itu”, “ke”, dan sebagainya tergantung pada bahasa dan konteks data.

Tabel 6. Tahap *Stopword Removal*

Sebelum <i>Stopword Removal</i>	Setelah <i>Stopword Removal</i>
lebih fleksibel saja kalau kuliah online apalagi dengan berbagai teknologi yang sangat canggih yang dapat digunakan untuk mendukung proses perkuliahan	fleksibel kuliah online teknologi canggih mendukung proses perkuliahan
pelajaran waktu berkuliah online se iyya kurang efektif saya rasa sebab terbatasnya kuota dan jaringan pada saat online ada waktu saat kuliah online juga jaringan kurang stabil kalau mati listrik jadi otomatis berdampak ke	pelajaran berkuliah online kurang efektif terbatasnya kuota jaringan online kuliah online jaringan kurang stabil mati listrik otomatis

jaringan dan interaksi ke dosen terbatas serta tidak dapat tau berinteraksi ke sesama teman	berdampak jaringan interaksi dosen terbatas tidak tau berinteraksi teman
karena kadang mahasiswa tidak memperhatikan saat dosen menjelaskan materinya dan sebagian juga mahasiswa lainnya dapat kendala seperti jaringan kurang bagus	kadang mahasiswa tidak memperhatikan dosen materinya mahasiswa kendala jaringan kurang bagus
tidak efektif	tidak efektif
karena kalau mata kuliahnya cuma itu nanggung kalau ke kampus	mata kuliahnya nanggung kampus

Tabel 6 menunjukkan hasil dari *stopword removal* yang sebelumnya merupakan hasil dari *tokenizing*. Proses ini mengurangi kata-kata yang tidak berguna. Contohnya, kata-kata yang ada pada opini kelima “karena”, “kalau”, “cuma”, “itu”, dan “ke” akan dihapus, yang menyisakan token-token seperti mata, kuliahnya, nanggung, kampus. Begitupun dengan opini lainnya.

3.2.6 Tokenizing

Tokenizing atau tokenisasi adalah proses untuk menguraikan teks atau kalimat kedalam bagian-bagian lebih kecil atau biasa disebut token. Token dapat berupa kata atau frasa.

Tabel 7. Tahap *Tokenizing*

Sebelum <i>Tokenizing</i>	Setelah <i>Tokenizing</i>
fleksibel kuliah online teknologi canggih mendukung proses perkuliahan	['fleksibel', 'kuliah', 'online', 'teknologi', 'canggih', 'mendukung', 'proses', 'perkuliahan']
pelajaran berkuliah online kurang efektif terbatasnya kuota jaringan online kuliah online jaringan kurang stabil mati listrik otomatis berdampak jaringan interaksi dosen terbatas tidak tau berinteraksi teman	['pelajaran', 'berkuliah', 'online', 'kurang', 'efektif', 'terbatasnya', 'kuota', 'jaringan', 'online', 'kuliah', 'online', 'jaringan', 'kurang', 'stabil', 'mati', 'listrik', 'otomatis', 'berdampak', 'jaringan', 'interaksi', 'dosen', 'terbatas', 'tidak', 'tau', 'berinteraksi', 'teman']
kadang mahasiswa tidak memperhatikan dosen materinya mahasiswa kendala jaringan kurang bagus	['kadang', 'mahasiswa', 'tidak', 'memperhatikan', 'dosen', 'materinya', 'mahasiswa', 'kendala', 'jaringan', 'kurang', 'bagus']
tidak efektif	['tidak', 'efektif']
mata kuliahnya nanggung kampus	['mata', 'kuliahnya', 'nanggung', 'kampus']

Pada tabel 7 terdapat hasil dari *tokenizing* yang didapatkan berdasarkan dari hasil *word correction* sebelumnya. Adapun pada opini kelima dipecah menjadi beberapa bagian, yaitu ['mata', 'kuliahnya', 'nanggung', 'kampus']. Begitupun dengan opini lainnya yang akan diuraikan.

3.2.7 Stemming

Tujuan proses *stemming*, yang merupakan bagian dari *preprocessing* text mining, adalah untuk mengkonversi kata ke bentuk dasar atau bentuk dasar (*stem*) atau membuat kata menjadi bentuk yang lebih sederhana. Pada proses ini digunakan hasil dari deteksi bahasa untuk membedakan *stemming* bahasa indonesia dan inggris.

Tabel 8. Tahap *Stemming*

Sebelum <i>Stemming</i>	Setelah <i>Stemming</i>
['fleksibel', 'kuliah', 'online', 'teknologi', 'canggih', 'mendukung', 'proses', 'perkuliahan']	['fleksibel', 'kuliah', 'online', 'teknologi', 'canggih', 'dukung', 'proses', 'kuliah']
['pelajaran', 'berkuliah', 'online', 'kurang', 'efektif', 'terbatasnya', 'kuota', 'jaringan', 'online', 'kuliah', 'online', 'jaringan', 'kurang', 'stabil', 'mati', 'listrik', 'otomatis', 'berdampak', 'jaringan', 'interaksi', 'dosen', 'terbatas', 'tidak', 'tau', 'berinteraksi', 'teman']	['ajar', 'kuliah', 'online', 'kurang', 'efektif', 'batas', 'kuota', 'jaring', 'online', 'kuliah', 'online', 'jaring', 'kurang', 'stabil', 'mati', 'listrik', 'otomatis', 'dampak', 'jaring', 'interaksi', 'dosen', 'batas', 'tidak', 'tau', 'interaksi', 'teman']
['kadang', 'mahasiswa', 'tidak', 'memperhatikan', 'dosen', 'materinya', 'mahasiswa', 'kendala', 'jaringan', 'kurang', 'bagus']	['kadang', 'mahasiswa', 'tidak', 'perhati', 'dosen', 'materi', 'mahasiswa', 'kendala', 'jaring', 'kurang', 'bagus']
['tidak', 'efektif']	['tidak', 'efektif']
['mata', 'kuliahnya', 'nanggung', 'kampus']	['mata', 'kuliah', 'nanggung', 'kampus']

Hasil dari *tokenizing* sebelumnya kemudian akan melalui proses *stemming* yang terdapat pada tabel 8 menunjukkan adanya proses konversi kata-kata ke bentuk dasarnya. Sebagai contoh, pada opini kelima, kata (kuliahnya) diubah menjadi (kuliah) dan begitupun dengan kata-kata yang mengalami perubahan serupa pada opini lainnya.

3.3 Pelabelan Data

Setelah menyelesaikan tahap pembersihan dalam proses *preprocessing* didapatkan dari 1680 data yang dikumpulkan tersisa 1.679 data yang bersih, kemudian tahap berikutnya yakni pelabelan data. Pada tahap ini, hasil dari *preprocessing* akan digolongkan ke dalam label yang sesuai. Ada dua label atau sentimen yang akan digunakan, yaitu positif dan negatif. Pelabelan data diproses menggunakan model BERT yakni: “nlptown/bert-base-multilingual-uncased-sentiment”. Model ini memberikan label sentimen dalam bentuk skor 1 hingga 5 bintang, yang kemudian dikonversi menjadi positif dan negatif.

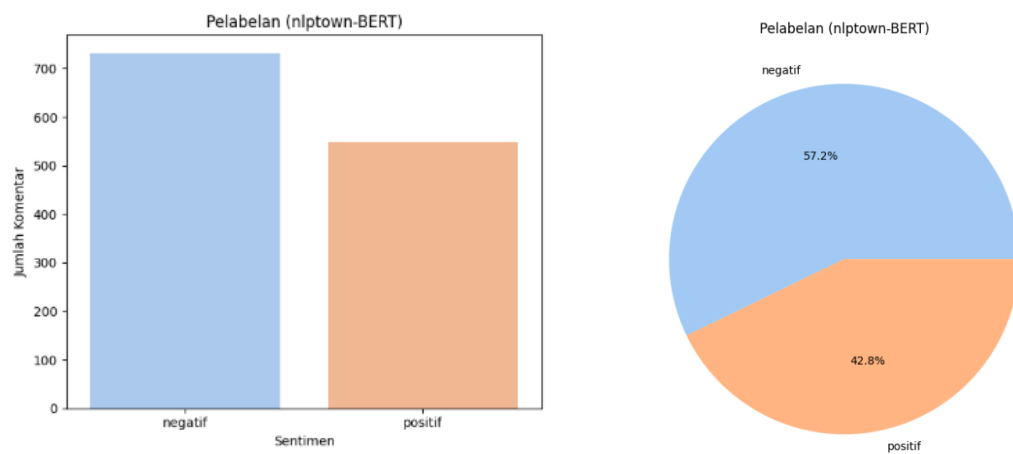
Tabel 9. Tahap Pelabelan Data

Hasil <i>Preprocessing</i>	Label Asli	Konversi
['fleksibel', 'kuliah', 'online', 'teknologi', 'canggih', 'dukung', 'proses', 'kuliah']	4 stars	Positif
['ajar', 'kuliah', 'online', 'kurang', 'efektif', 'batas', 'kuota', 'jaring', 'online', 'kuliah', 'online', 'jaring', 'kurang', 'stabil', 'mati', 'listrik', 'otomatis', 'dampak', 'jaring', 'interaksi', 'dosen', 'batas', 'tidak', 'tau', 'interaksi', 'teman']	2 stars	Negatif

['kadang', 'mahasiswa', 'tidak', 'perhati', 'dosen', 'materi', 'mahasiswa', 'kendala', 'jaring', 'kurang', 'bagus']	2 stars	Negatif
['tidak', 'efektif']	1 stars	Negatif
['mata', 'kuliah', 'nanggung', 'kampus']	5 stars	Positif

Tabel 9 merupakan hasil dari tahap *preprocessing* sebelumnya dan hasil dari pelabelan data. Proses ini membagi kategori komentar ke dalam label sentimen seperti positif dan negatif. Pelabelan dilakukan dengan mengklasifikasikan setiap teks ke dalam lima kategori sentimen berdasarkan jumlah bintang, mulai dari 1 hingga 5 bintang, yang merepresentasikan tingkat kepuasan atau sentimen dari sangat negatif hingga sangat positif. Hasil pelabelan kemudian dikonversi menjadi dua kelas, di mana label 1–2 bintang dikategorikan sebagai sentimen negatif, label 4–5 bintang dimasukkan ke dalam kategori positif, kemudian untuk label 3 bintang yang biasanya dianggap netral akan dihilangkan. Dalam banyak penelitian yang berfokus pada analisis sentimen, pengabaian kelas netral telah menjadi praktik yang umum dan diakui.

Pendekatan ini sering didasari oleh asumsi bahwa ulasan atau teks dengan polaritas netral memiliki nilai informatif yang terbatas dalam membedakan antara sentimen positif dan negatif yang jelas. Misalnya, dalam penelitian seperti yang diacu dalam [23], [24], [25] kelas netral sering kali dibuang untuk menyederhanakan masalah klasifikasi menjadi klasifikasi biner, yang berpotensi meningkatkan akurasi untuk identifikasi sentimen ekstrem. Pengabaian ini juga dapat dijelaskan oleh kenyataan bahwa sentimen netral sering kali ambigu atau kurang relevan untuk tujuan praktis, seperti ketika fokus utamanya adalah mengidentifikasi kepuasan atau ketidakpuasan pelanggan secara eksplisit, bukan pada opini yang tidak memiliki kecenderungan kuat.

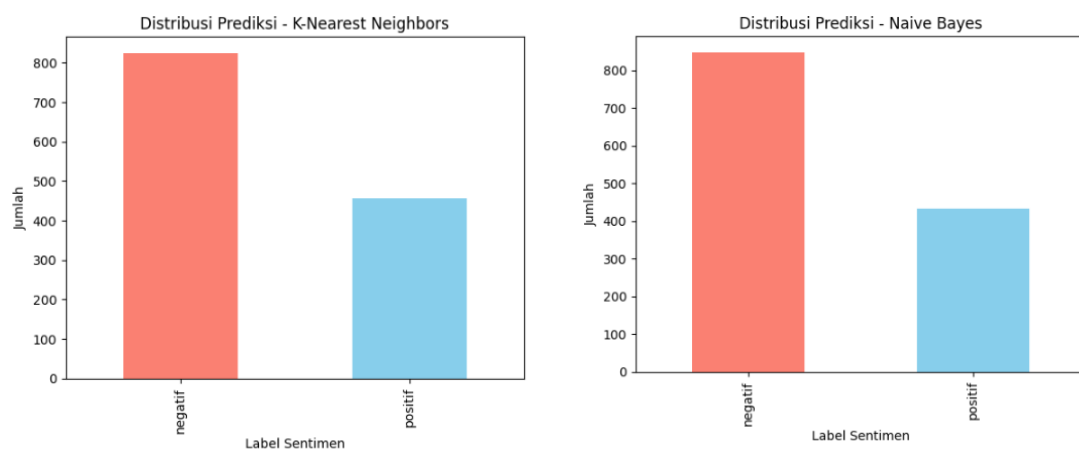


Gambar 2. Hasil Pelabelan

Gambar 2 merupakan tampilan dari hasil pelabelan sebelumnya yang menunjukkan bahwa hasil pelabelan lumayan seimbang serta adanya pengurangan jumlah opini karena penghapusan label netral. Pada tahapan *preprocessing* jumlah data berjumlah 1.679 opini, tetapi setelah pelabelan dan menghapus opini netral data berkurang menjadi 1.280 opini, dengan sentimen positif yang mencapai 42.8% yakni 548 opini dan sentimen negatif sebanyak 57.2% yakni 732 opini. Setelah pelabelan data dilakukan, langkah selanjutnya adalah proses pembobotan TF-IDF. Pemberian bobot menggunakan metode TF-IDF akan menentukan bobot setiap kata berdasarkan frekuensi kemunculannya, sehingga kosakata yang lebih sering muncul dalam suatu dokumen namun memiliki frekuensi rendah dalam keseluruhan dataset akan dianggap lebih penting.

3.4 Klasifikasi

Setelah melalui tahap *preprocessing*, pelabelan data, dan pembobotan TF-IDF, maka beralih ke tahap klasifikasi. Teknik klasifikasi yang digunakan dalam penelitian analisis sentimen ini menerapkan metode *K-Nearest Neighbor* (KNN) dan *Naïve Bayes*. Untuk proses evaluasi kinerja model, digunakan teknik *k-fold cross-validation* yang membagi data ke dalam beberapa lipatan secara stratifikasi sehingga setiap data diberi peluang yang setara untuk digunakan sebagai data pelatihan maupun data pengujian. Klasifikasi dilakukan untuk melihat seberapa banyak opini yang menerima ataupun menolak kebijakan perkuliahan *online* tersebut. Adapun algoritma yang digunakan adalah algoritma *K-Nearest Neighbor* dan *Naïve Bayes*. Nilai k pada algoritma KNN ditentukan secara empiris dengan menguji rentang nilai dari 1 hingga 20 menggunakan evaluasi *10-fold cross-validation*, selaras dengan praktik umum dalam literatur dan tutorial algoritma. Pengujian tersebut menghasilkan $k = 19$ sebagai nilai optimal karena memberikan akurasi rata-rata tertinggi.



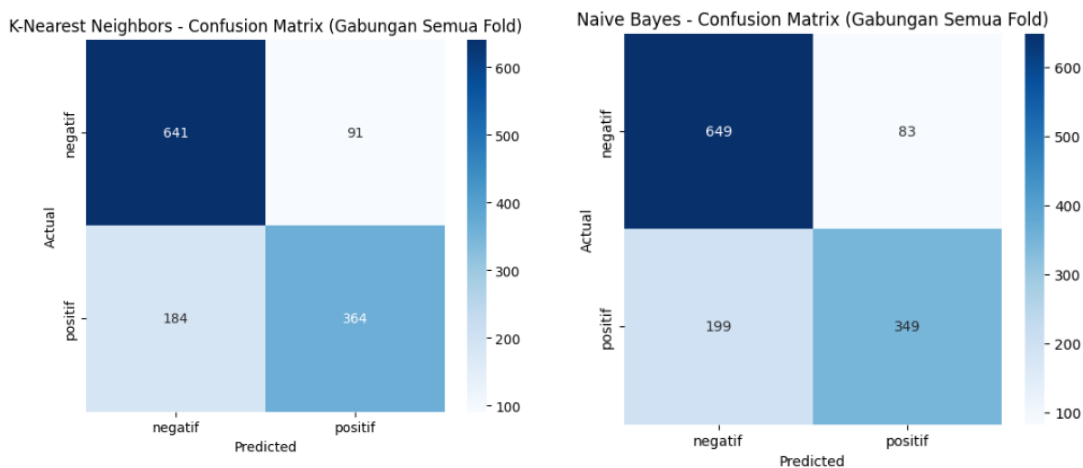
Gambar 3. Hasil Klasifikasi Berdasarkan Algoritma KNN dan *Naïve Bayes*

Gambar 3 merupakan hasil klasifikasi yang sudah dilakukan dan terdapat masing-masing hasil dari kedua algoritma, yakni dari 1280 dataset keseluruhan pada algoritma *K-Nearest Neighbor* terdapat 455 kategori positif dan 825 kategori negatif. Kemudian pada algoritma *Naïve Bayes* terdapat 432 kategori positif dan 848 kategori negatif.

3.5 Evaluasi

Langkah terakhir adalah evaluasi model yang sudah dibuat. Performa model dinilai berdasarkan perhitungan *accuracy*, *precision*, *recall*, dan *f-1 score*. Evaluasi ini dilakukan guna mengukur seberapa efektif kinerja model klasifikasi KNN dan *Naïve Bayes* dalam memprediksi data secara akurat serta memberikan gambaran keseluruhan tentang kinerja model tersebut. Untuk menghindari overfitting dan memastikan performa model tetap stabil saat diterapkan pada data baru, maka digunakan teknik *k-fold cross validation*. Teknik ini mendistribusikan data menjadi sejumlah *fold* atau subset, kemudian model dilatih dan diuji secara bergantian pada setiap *fold*, sehingga hasil evaluasi menjadi lebih stabil dan representatif.

Hasil pengujian akan divisualisasikan dalam *confusion matrix*. Matriks ini dapat dimanfaatkan untuk mengukur metrik evaluasi seperti *accuracy*, *precision*, *recall*, dan *f-1 score* pada setiap kelas sentimen. Matriks ini akan berperan dalam menentukan tingkat akurasi model klasifikasi mengenali tiap kelas dan frekuensi terjadinya kesalahan antar kelas terjadi. Matriks ini akan menyajikan ilustrasi yang lebih mendalam mengenai cara kerja model klasifikasi bertugas untuk menentukan label sentimen dengan tepat.



Gambar 4. *Confusion Matrix* Algoritma KNN dan *Naïve Bayes*

Gambar 4 menyajikan tampilan dari visualisasi *confusion matrix* pada kedua algoritma. Evaluasi performa model *K-Nearest Neighbor* dan *Naïve Bayes* dalam klasifikasi sentimen diterapkan melalui *confusion matrix*. Tabel 10 menyajikan ringkasan hasil pengujian tersebut:

Tabel 10. Hasil Evaluasi

Metode	Model	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
Setelah <i>K-Fold Cross Validation</i>	<i>K-Nearest Neighbor</i>	78.52%	78.72%	78.52%	78.14%
	<i>Naïve Bayes</i>	77.97%	78.46%	77.97%	77.45%

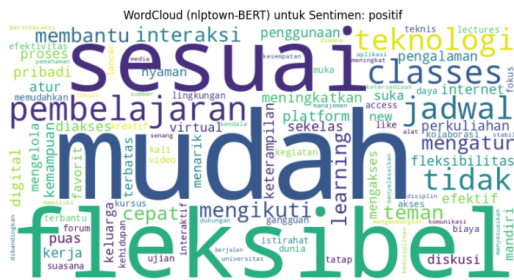
Tabel 10 menyajikan hasil evaluasi kinerja dari dua algoritma klasifikasi, yaitu *K-Nearest Neighbor* (KNN) dan *Naïve Bayes*, setelah dilakukan proses *K-Fold Cross Validation*. Tujuan dari penerapan metode ini dimaksudkan untuk memperoleh estimasi kinerja model yang lebih stabil dan menyeluruh, dengan mengurangi risiko *overfitting* akibat pembagian data yang tidak representatif. Berdasarkan hasil yang ditampilkan dengan menggunakan nilai $k = 10$, algoritma KNN menunjukkan performa yang sedikit lebih unggul daripada *Naïve Bayes* pada semua metrik evaluasi. KNN menghasilkan *accuracy* sebesar 78.52%, *precision* 78.72%, *recall* 78.52%, dan *f1-score* 78.14%. Sementara itu, *Naïve Bayes* mencatatkan nilai *accuracy* sebesar 77.97%, dengan *precision* 78.46%, *recall* 77.97%, dan *f1-score* 77.45%.

Secara keseluruhan, algoritma KNN menunjukkan performa yang paling optimal setelah melalui evaluasi menggunakan *K-Fold Cross Validation*, dengan pencapaian metrik yang melampaui *Naïve Bayes* pada semua aspek evaluasi. Meskipun demikian, selisih performa antara keduanya tergolong kecil, yang menunjukkan bahwa *Naïve Bayes* tetap mampu memberikan hasil yang kompetitif dan stabil. Hasil ini menegaskan bahwa pemilihan metode evaluasi seperti *K-Fold Cross Validation* sangat penting untuk mendapatkan gambaran kinerja model yang lebih realistis, serta menekankan pentingnya menyesuaikan algoritma dengan karakteristik data latih yang tersedia.

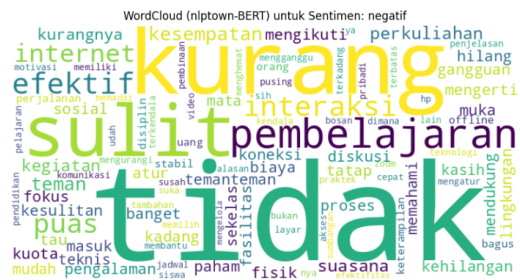
3.6 Visualisasi

Word cloud adalah salah satu bentuk visualisasi untuk memvisualisasikan kata dengan frekuensi kemunculan tertinggi dengan skala yang lebih besar, sedangkan kata-kata dengan

frekuensi rendah tetap akan ditampilkan, namun dengan skala yang lebih kecil. Berikut adalah gambar visualisasi opini dengan sentimen positif dan negatif



Gambar 5. Word Cloud Sentimen Positif



Gambar 6. Word Cloud Sentimen Negatif

Gambar 5 menunjukkan kata-kata dari opini positif seperti *membantu*, *mudah*, *fleksibel*, *puas*, *efektif*, dan lainnya yang mencerminkan pandangan positif terhadap kemudahan dan dukungan dalam perkuliahan *online*. Adapun gambar 6 memuat kata-kata seperti *sulit*, *kurang*, *tidak*, *gangguan* dan lainnya yang mengindikasikan kendala dan ketidakpuasan mahasiswa terhadap pelaksanaan kuliah daring.

4. KESIMPULAN

Berdasarkan hasil yang didapatkan, penerapan algoritma *K-Nearest Neighbor* (KNN) dan *Naïve Bayes* dalam analisis sentimen terhadap kebijakan perkuliahan *online* multibahasa menunjukkan kinerja yang tergolong cukup baik, terkhusus pada algoritma *K-Nearest Neighbor* (KNN). Salah satu aspek pembeda pada penelitian ini mencakup penggunaan data yang tidak hanya berasal dari bahasa Indonesia, tetapi juga mencakup bahasa lokal, yaitu bahasa Mandar. Hal ini menambah tingkat kompleksitas, khususnya dalam tahap *preprocessing* seperti normalisasi, penghapusan *stopword*, dan *stemming*, mengingat keterbatasan sumber daya linguistik untuk bahasa Mandar. Dataset diperoleh melalui penyebaran kuesioner yang menampung opini mahasiswa mengenai kebijakan perkuliahan *online* dalam berbagai bahasa, dengan total sebanyak 1.680 data. Setelah *preprocessing* yang dilakukan, dataset hasil *preprocessing* berjumlah 1679 entri. Pada proses pelabelan juga terdapat pengurangan data, karena opini yang termasuk label netral (3 bintang) akan dihapus yang menghasilkan sebanyak 1280 dataset keseluruhan. Kemudian klasifikasi akan dilakukan untuk masing-masing algoritma, hasil yang didapatkan dari algoritma *K-Nearest Neighbor* yakni 455 kategori positif dan 825 kategori negatif serta pada algoritma *Naïve Bayes* terdapat 432 kategori positif dan 848 kategori negatif.

Evaluasi dilakukan menggunakan metode *confusion matrix* dan *k-fold cross validation* untuk menghindari *overfitting*. Hasil evaluasi menunjukkan bahwa KNN menghasilkan *accuracy* sebesar 78.52%, *precision* 78.72%, *recall* 78.52%, dan *f1-score* 78.14%. Sementara itu, *Naïve Bayes* mencatatkan nilai *accuracy* sebesar 77.97%, dengan *precision* 78.46%, *recall* 77.97%, dan *f1-score* 77.45%. Hal ini membuktikan bahwa algoritma *K-Nearest Neighbor* memiliki kemampuan generalisasi yang sedikit memiliki performa yang lebih tinggi daripada *Naïve Bayes* dalam mengklasifikasikan sentimen mahasiswa terhadap kebijakan perkuliahan *online*, terutama dalam konteks data multibahasa yang memiliki kompleksitas tinggi.

Selain itu, penggunaan *k-fold cross validation* menyajikan ilustrasi performa model yang lebih akurat dan stabil karena seluruh data dilibatkan dalam proses pelatihan dan pengujian secara bergantian. Hal ini penting dalam konteks dataset yang heterogen dan terbatas seperti ini, di mana distribusi opini dan bahasa yang bervariasi dapat memengaruhi kualitas prediksi model. Oleh karena itu, hasil penelitian ini tidak terbatas pada pemberian wawasan terkait dengan seberapa

efektif algoritma klasifikasi dalam analisis sentimen multibahasa, tetapi juga menegaskan pentingnya pemilihan metode evaluasi dan strategi praproses data yang sesuai. Penelitian ini diharapkan bisa dijadikan sebagai acuan dalam pengembangan sistem analisis sentimen yang lebih inklusif terhadap keragaman bahasa lokal di masa mendatang. Untuk penelitian selanjutnya, disarankan agar proses pelabelan ditingkatkan, terutama mengingat adanya penggunaan bahasa daerah seperti bahasa Mandar yang menyulitkan proses anotasi. Pendekatan pelabelan manual berbasis kamus lexicon untuk bahasa Mandar dapat menjadi solusi dalam meningkatkan akurasi label atau digunakan pada penelitian selanjutnya yang membahas mengenai multibahasa. Selain itu, pengembangan lebih lanjut dapat dilakukan melalui optimasi parameter, perluasan jumlah data latih, serta eksplorasi teknik lanjutan seperti feature selection atau penerapan algoritma yang lebih kompleks, seperti *Support Vector Machine* (SVM), *Random Forest*, *Convolutional Neural Network* (CNN), maupun pendekatan *deep learning* lainnya.

DAFTAR PUSTAKA

- [1] T. W. Putra, A. Triayudi, and Andrianingsih, "Analisis Sentimen Pembelajaran Daring menggunakan Metode Naïve Bayes, KNN, dan Decision Tree," *Jurnal Teknologi Informasi dan Komunikasi*, vol. 6, no. 1, pp. 20–26, 2022, doi: 10.35870/jti.
- [2] Samsir, Ambiyar, U. Verawardina, F. Edi, and R. Watrianthos, "Analisis Sentimen Pembelajaran Daring Pada Twitter di Masa Pandemi COVID-19 Menggunakan Metode Naïve Bayes," *JURNAL MEDIA INFORMATIKA BUDIDARMA*, vol. 5, no. 1, pp. 157–163, Jan. 2021, doi: 10.30865/mib.v5i1.2604.
- [3] N. Eldha Oktaviana and Y. Arum Sari, "Analisis Sentimen Terhadap Kebijakan Kuliah Daring Selama Pandemi Menggunakan Pendekatan Lexicon Based Features dan Support Vector Machine," *Jurnal Teknologi Informasi dan Ilmu Komputer*, 2022, doi: 10.25126/jtiik.202295625.
- [4] Zulfachmi, N. M. Kautsar, Z. A. Ramadhan, M. B. Mbembe, and Yvonme, "Analisis Sentimen Tanggapan Masyarakat Terhadap Pembelajaran Daring di Masa Pandemi COVID-19 Melalui Sosial Media Twitter Menggunakan Klasifikasi Naïve Bayes," *Prosiding Seminar Implementasi Teknologi Informasi dan Komunikasi*, vol. 2, no. 1, 2023, doi: 10.31284/p.semtik.2023-1.3938.
- [5] A. A. P. T. Mara, E. Sedyono, and H. Purnomo, "Penerapan Algoritma K-Nearest Neighbors Pada Analisis Sentimen Metode Pembelajaran Dalam Jaringan (DARING) Di Universitas Kristen Wira Wacana Sumba," Jun. 2021.
- [6] Amelia, "Teks dan Analisis Sentimen Pada Chat Grup Whatsapp Menggunakan Long Short Term Memory (LSTM)," vol. 3, pp. 1–23, Jan. 2023.
- [7] O. Mailani, I. Nuraeni, S. A. Syakila, and J. Lazuardi, "Bahasa Sebagai Alat Komunikasi Dalam Kehidupan Manusia," Online, Jan. 2022. [Online]. Available: www.plus62.isha.or.id/index.php/kampret
- [8] N. Widya Utami and M. Artana, "Text Mining Dalam Analisis Sentimen Pembelajaran Daring di Masa Pandemi Covid 19 Menggunakan Algoritma K-Nearest Neighbor," 2022.
- [9] S. N. Z. Wati, H. Herman, and H. Darwis, "Naive Bayes Classifier dan K-Nearest Neighbor pada Analisis Sentimen Perkuliahan Daring di Universitas Muslim Indonesia," *Buletin Sistem Informasi dan Teknologi Islam*, vol. 5, no. 1, pp. 47–54, Apr. 2024, doi: 10.33096/busiti.v5i1.2202.
- [10] E. M. A. Ernamia and A. Herliana, "Analisis Sentimen Kuliah Daring dengan Algoritma Naïve Bayes, K-NN dan Decision Tree," *JURNAL RESPONSIF*, vol. 4, no. 1, pp. 70–80, 2022, [Online]. Available: <https://ejurnal.ars.ac.id/index.php/jti>
- [11] W. Khofifah, D. N. Rahayu, and A. M. Yusuf, "Analisis Sentimen Menggunakan Naive Bayes Untuk Melihat Review Masyarakat Terhadap Tempat Wisata Pantai Di Kabupaten Karawang Pada Ulasan Google Maps," *Jurnal Interkom: Jurnal Publikasi Ilmiah Bidang Teknologi Informasi dan Komunikasi*, vol. 16, no. 4, pp. 28–38, Jan. 2022, doi: 10.35969/interkom.v16i4.192.
- [12] A. Apriani, H. Zakiyudin, and K. Marzuki, "Penerapan Algoritma Cosine Similarity dan Pembobotan TF-IDF System Penerimaan Mahasiswa Baru pada Kampus Swasta," *Jurnal Bumigora Information Technology (BITE)*, vol. 3, no. 1, pp. 19–27, Jul. 2021, doi: 10.30812/bite.v3i1.1110.

- [13] M. F. Naufal and S. F. Kusuma, "Analisis Sentimen pada Media Sosial Twitter Terhadap Kebijakan Pemberlakuan Pembatasan Kegiatan Masyarakat Berbasis Deep Learning," *JEPIN(Jurnal Edukasi dan Penelitian Informatika)*, vol. 8, pp. 44–49, Apr. 2022.
- [14] W. Irmayani, "Persepsi Publik Terhadap Kenaikan PPN 12%: Pendekatan Sentimen Pada Komentar Youtube," *JURNAL KHATULISTIWA INFORMATIKA*, vol. 12, no. 2, pp. 112–8, Dec. 2024.
- [15] K. Mustaqim, F. A. Maresti, and I. N. Dewi, "Analisis Sentimen Ulasan Aplikasi PosPay untuk Meningkatkan Kepuasan Pengguna dengan Metode K-Nearest Neighbor (KNN)," *Edumatic: Jurnal Pendidikan Informatika*, vol. 8, no. 1, pp. 11–20, Jun. 2024, doi: 10.29408/edumatic.v8i1.24779.
- [16] M. Furqan, Sriani, and S. M. Sari, "Analisis Sentimen Menggunakan K-Nearest Neighbor Terhadap New Normal Masa Covid-19 Di Indonesia," Feb. 2022. [Online]. Available: www.tripadvisor.com
- [17] C. P. A. Mulya, P. Nugraha, and I. Santoso, "Analisis Sentimen Masyarakat Terhadap Pembangunan Kereta Cepat Jakarta-Bandung Menggunakan Algoritma K-Nearest Neighbors (KNN)," Jul. 2023. [Online]. Available: www.gataframework.com/textmining. Peneliti
- [18] Dewi Rezekika, "Penerapan Algoritma Naïve Bayes Untuk Memprediksi Penjualan Spare Part Sepeda Motor," *Jurnal Pelita Informatika*, vol. 8, 2020.
- [19] I. K. A. Sugiarta, P. A. C. Dewi, and N. W. Utami, "Analisa Sentimen Mahasiswa Terhadap Layanan STMIK Primakara Menggunakan Algoritma Naive Bayes dan K-Nearest Neighbor," Aug. 2023.
- [20] H. Santoso, Armansyah, and D. Desliani, "Analisis Sentimen Mahasiswa Terkait Pembelajaran Tatap Muka Menggunakan Metode Naïve Bayes Classifier," Aug. 2022.
- [21] M. H. Casandy and D. Mahdiana, "Penerapan Algoritma Naïve Bayes Untuk Melakukan Analisis Sentimen Pada PT Pos Indonesia (Persero) Application of Nave Bayes Algorithm to Perform Sentiment Analysis at PT Pos Indonesia (Persero)," *KRESNA: Jurnal Riset dan Pengabdian Masyarakat*, vol. 2, no. 2, pp. 213–221, 2022, [Online]. Available: <https://jurnaldrpm.budiluhur.ac.id/index.php/Kresna/>
- [22] A. Muthalib, *KAMUS BAHASA MANDAR-INDONESIA*. UjungPandang, 1977.
- [23] S. Thomas, Yuliana, and N. P., "Studi Analisis Metode Analisis Sentimen pada YouTube," *JIFOTECH (JOURNAL OF INFORMATION TECHNOLOGY)*, vol. 1, no. 1, pp. 1–7, Mar. 2021.
- [24] G. Nkhata, U. Anjum, and J. Zhan, "Sentiment Analysis of Movie Reviews Using BERT," Feb. 2025, [Online]. Available: <http://arxiv.org/abs/2502.18841>
- [25] R. Socher *et al.*, "Recursive Deep Models for Semantic Compositionality Over a Sentiment Treebank," Association for Computational Linguistics, 2013. [Online]. Available: <http://nlp.stanford.edu/>