

BAB 4

HASIL PENELITIAN

4.1 RINGKASAN HASIL PENELITIAN

Berdasarkan penelitian tentang analisis sentimen komentar pengguna TikTok terhadap penggunaan ChatGPT dalam skripsi menggunakan metode Naïve Bayes. Data diperoleh melalui proses scraping menggunakan tools Tapicker, yang menghasilkan sebanyak 1.200 komentar dari beberapa postingan TikTok yang membahas topik terkait. Sebelum analisis lebih lanjut, data komentar diproses terlebih dahulu. Proses ini meliputi penghapusan huruf besar, pembersihan, normalisasi, tokenisasi, penghapusan kata yang tidak memiliki makna penting, dan stemming. Setelah diproses terlebih dahulu, setiap komentar diklasifikasikan secara manual ke dalam tiga kategori yaitu sentimen positif, negatif, dan netral.

Dataset yang telah dibersihkan kemudian dibagi menjadi dua bagian, yakni 960 data untuk pelatihan (80%) dan 240 data untuk pengujian (20%). Pada data latih, terdapat 436 komentar bersentimen positif, 272 komentar negatif, dan 252 komentar netral. Sementara itu, data uji terdiri atas 123 komentar positif, 64 komentar negatif, dan 53 komentar netral.

4.2 AKURASI ALGORITMA

Pada tahap ini, kinerja algoritma Naïve Bayes dalam mengklasifikasikan komentar pengguna TikTok ke dalam tiga kategori sentimen dievaluasi. Proses klasifikasi melibatkan 240 set data uji yang sebelumnya diberi label secara manual. Berbagai metrik evaluasi digunakan untuk menilai kinerja model secara keseluruhan, akurasi, presisi, *recall*, dan skor F1. Metrik ini mengukur kemampuan model untuk mengidentifikasi sentimen secara akurat dan konsisten.

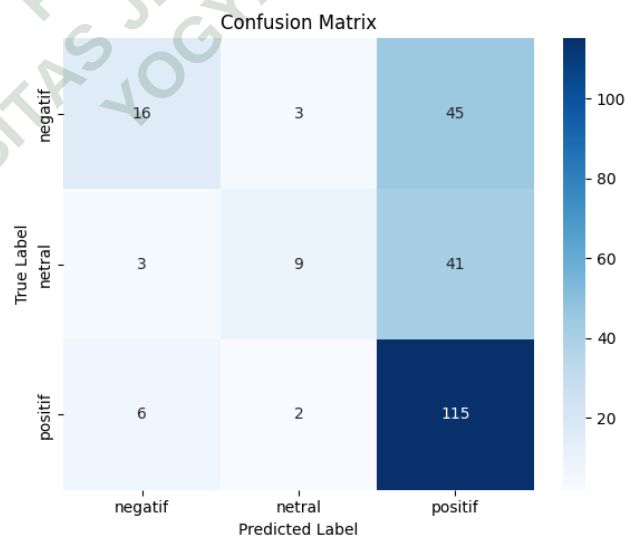
Hasil evaluasi menunjukkan bahwa akurasi model mencapai 58,33%. Detail hasil evaluasi performa model berdasarkan masing-masing kategori sentimen dapat dilihat pada Tabel 4.1.

Tabel 4.1 Hasil Precision, Recall, F1-score, dan Support

	Precision	Recall	F1-score	Support
Negatif	0.64	0.25	0.36	64
Netral	0.64	0.17	0.27	53
Positif	0.57	0.93	0.71	123

Berdasarkan Tabel 4.1 dapat disimpulkan bahwa algoritma Naïve Bayes paling optimal dalam mengklasifikasikan komentar positif, yang ditunjukkan oleh nilai *recall* sebesar 0,93 dan *f1-score* tertinggi sebesar 0,71. Sebaliknya, performa klasifikasi untuk sentimen negatif dan netral masih rendah, dengan nilai *recall* masing-masing hanya 0,25 dan 0,17.

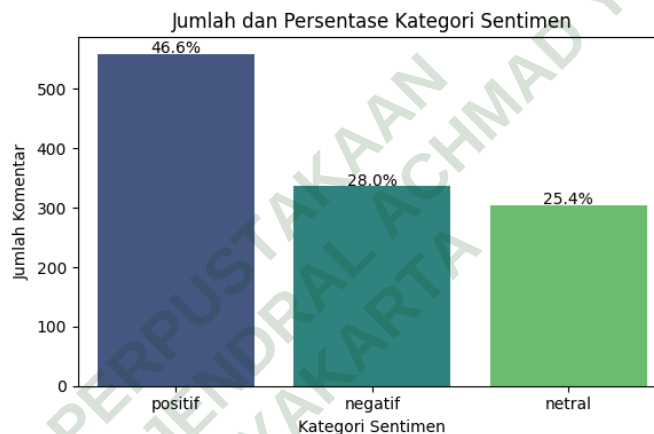
Hal ini mengindikasikan bahwa model cenderung mengklasifikasikan komentar ke dalam kategori positif, dan masih mengalami kesulitan dalam mengenali komentar negatif dan netral. Ketidakseimbangan data dan fitur representasi yang kurang kuat dapat menjadi salah satu penyebab rendahnya performa pada dua kategori tersebut. Kemudian terdapat hasil *confusion matrix* yang dapat dilihat pada Gambar 4.1.

**Gambar 4.1** *Confusion Matrix*

Pada Gambar 4.1 menunjukkan bahwa model lebih akurat dalam mengklasifikasikan komentar positif, namun masih sering keliru dalam mengklasifikasi komentar negatif, dan netral.

4.3 HASIL ANALISIS

Hasil klasifikasi sentimen dibagi menjadi tiga kategori, yaitu sentimen positif, negatif, dan netral. Proses klasifikasi menghasilkan distribusi sentimen berdasarkan total data komentar yang dianalisis. Distribusi ini digambarkan secara grafik pada Gambar 4.2.



Gambar 4.2 Grafik Presentase Kategori Sentimen

Pada Gambar 4.2 dapat dilihat bahwa komentar dengan kategori sentimen positif mendominasi sebanyak 47%, diikuti oleh komentar negatif sebanyak 28% dan komentar netral sebesar 25%. Kategori komentar positif mencerminkan tanggapan yang bersifat mendukung, seperti ChatGPT sangat membantu dalam penyusunan skripsi. Berikut contoh data komentar dengan sentimen positif dapat dilihat pada Tabel 4.2.

Tabel 4.2 Komentar Positif

No	Komentar Positif
1.	pkony chatgpt trbaik dah ngbntu bgttt
2.	ttp bijak menggunakan chatgpt
3.	bagus chat gpt . terbaik
4.	beneran ngebantu banget!

5.	chatgpt membuat proses penulisan skripsi jadi lebih cepat
----	---

Sementara itu, komentar negatif menunjukkan kekhawatiran terhadap penggunaan ChatGPT yang dapat menurunkan kemampuan berpikir kritis dan menimbulkan masalah etika akademik seperti plagiarisme, serta menyajikan referensi yang terkadang kurang relevan atau tidak akurat. Berikut contoh komentar negatif dapat dilihat pada Tabel 4.3.

Tabel 4.3 Komentar Negatif

No	Komentar Negatif
1.	tapi kutipan dari referensinya kurang sesuai dan ada yg ga sesuai juga
2.	gpt bikin orang males mikir
3.	skripsi harus karya pemikiran sendiri, bukan karya mesin
4.	chat gpt kalau diturnitin kena plagiaris
5.	kejahatan akademik

Sedangkan komentar netral umumnya berisi tanggapan yang bersifat informatif, pertanyaan, atau tidak memihak komentar positif maupun komentar negatif. Berikut contoh data komentar netral dapat dilihat pada Tabel 4.4.

Tabel 4.4 Komentar Netral

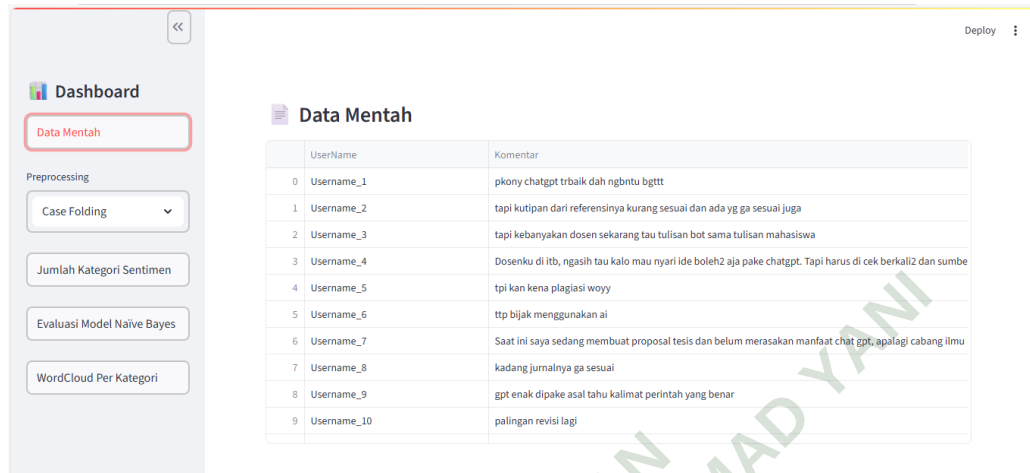
No	Komentar Netral
1.	apakah tidak berbahaya
2.	bisa langsung dari bab 1-5 ga? 😊
3.	gk ketahuan kah
4.	itu sumber referensi ai nya akurat ga?
5.	pakai gpt tapi ketik ulang itu ketahuan gasih 😊

4.4 VISUALISASI DATA

4.4.1 Tampilan Data Mentah

Tampilan data mentah digunakan untuk menampilkan data yang belum dilakukan tahap *preprocessing*. Data ini masih dalam bentuk asli sebagaimana

diambil dari platform TikTok. Tampilan data mentah tersebut dapat dilihat pada Gambar 4.3.



	Username	Komentar
0	Username_1	pkony chatgpt trbaik dah ngbntu bgttt
1	Username_2	tapi kutipan dari referensinya kurang sesuai dan ada yg ga sesuai juga
2	Username_3	tapi kebanyakan dosen sekarang tau tulisan bot sama tulisan mahasiswa
3	Username_4	Dosenku di itb, ngasih tau kalo mau nyari ide boleh2 aja pake chatgpt. Tapi harus di cek berkaliz dan sumbe
4	Username_5	tpi kan kena plagiasi woyy
5	Username_6	ttp bijak menggunakan ai
6	Username_7	Saat ini saya sedang membuat proposal tesis dan belum merasakan manfaat chat gpt, apalagi cabang ilmu
7	Username_8	kadang jurnalnya ga sesuai
8	Username_9	gpt enak dipake asal tahu kalimat perintah yang benar
9	Username_10	palingan revisi lagi

Gambar 4.3 Tampilan Data Mentah

Pada Gambar 4.3 ditampilkan contoh visualisasi data mentah yang akan menjadi dasar untuk proses *cleaning*, *case folding*, normalisasi, *tokenizing*, *stopword removal*, dan *stemming*. Berikut kode untuk menampilkan data mentah ke *dashboard*.

```
df_raw = pd.read_excel("data_komentar.xlsx")
st.subheader("📄 Data Mentah")
st.dataframe(df_raw, use_container_width=True)
st.stop()
```

Kode tersebut membaca file `data_komentar.xlsx`, menampilkannya dalam bentuk tabel pada *dashboard* Streamlit, dan menghentikan proses eksekusi setelah data ditampilkan.

4.4.2 Tampilan Preprocessing

Untuk menampilkan hasil dari setiap tahapan preprocessing pada data komentar TikTok, berikut kode yang digunakan.

```
if not show_raw and not show_kategori and not show_evaluasi
and not show_piechart:
    filename = file_mapping[show_preprocessing]
    hasil_kolom = hasil_mapping[show_preprocessing]
    df = load_data(filename)
```

```

st.title(f"Preprocessing: {show_preprocessing}")
if df is not None:
    username_col = next((col for col in df.columns if
col.lower() == "username"), None)
    hasil_col = next((col for col in df.columns if
col.lower() == hasil_kolom.lower()), None)
    if username_col and hasil_col:
        st.dataframe( df[[username_col,
hasil_col]].rename(columns={
            username_col: "Username",
            hasil_col: f"Hasil {show_preprocessing}"
        })),
        use_container_width=True)

```

Dengan kode tersebut, dashboard dapat secara otomatis menampilkan hasil dari setiap tahapan *preprocessing*. Adapun tampilan hasilnya adalah sebagai berikut:

a. Tampilan Case Folding

Tampilan *case folding* menampilkan data komentar yang telah melalui proses case folding. Proses ini mengubah semua huruf dalam kalimat menjadi huruf kecil. Tampilan hasil dari proses case folding pada Gambar 4.4.



The screenshot shows a web dashboard with a sidebar on the left containing buttons for 'Data Mentah', 'Preprocessing' (with a dropdown menu showing 'Case Folding'), 'Jumlah Kategori Sentimen', 'Evaluasi Model Naïve Bayes', and 'WordCloud Per Kategori'. The main area is titled 'Preprocessing: Case Folding' and displays a table with two columns: 'Username' and 'Hasil Case Folding'. The table contains 10 rows of data, each showing a username and a corresponding comment that has been converted to all lowercase letters.

	Username	Hasil Case Folding
0	Username_1	pkony chatgpt trbaik dah ngbntu bgttt
1	Username_2	tapi kutipan dari referensinya kurang sesuai dan ada yg ga sesuai juga
2	Username_3	tapi kebanyakan dosen sekarang tau tulisan bot sama tulisan mahasiswa
3	Username_4	dosenku di itb, ngasih tau kalo mau nyari ide boleh2 aja pake chatgpt, tapi harus di cek berkali2 dan sumber
4	Username_5	tpi kan kena plagiasi woyy
5	Username_6	ttp bijak menggunakan ai
6	Username_7	saat ini saya sedang membuat proposal tesis dan belum merasakan manfaat chat gpt, apalagi cabang ilmu i
7	Username_8	kadang jurnalnya ga sesuai
8	Username_9	gpt enak dipake asal tahu kalimat perintah yang benar
9	Username_10	palingan revisi lagi

Gambar 4.4 Tampilan Case Folding

Pada Gambar 4.4 ditampilkan hasil proses *case folding*, di mana seluruh karakter dalam komentar telah diubah menjadi huruf kecil. Proses ini bertujuan untuk menyeragamkan teks sehingga kata-kata dianggap sebagai entitas yang

sama. Setelah tahap ini, data akan dilanjutkan ke proses cleaning untuk menghilangkan karakter yang tidak relevan.

b. Tampilan Cleaning

Tampilan cleaning menampilkan hasil komentar yang telah melalui proses pembersihan teks, seperti menghapus karakter khusus (simbol, angka, tanda baca), emotikon, serta tautan yang tidak relevan. Hasil dari proses ini dapat dilihat pada Gambar 4.5.



	Username	Hasil Cleaning
0	Username_1	pkony chatgpt trbaik dah ngntu bgttt
1	Username_2	tapi kutipan dari referensinya kurang sesuai dan ada yg ga sesuai juga
2	Username_3	tapi kebanyakan dosen sekarang tau tulisan bot sama tulisan mahasiswa
3	Username_4	dosenku di itb ngasih tau kalo mau nyari ide boleh aja pake chatgpt tapi harus di cek berkali dan sumberny .
4	Username_5	tpi kan kena plagiasi woyy
5	Username_6	ttp bijak menggunakan ai
6	Username_7	saat ini saya sedang membuat proposal tesis dan belum merasakan manfaat chat gpt apalagi cabang ilmu k
7	Username_8	kadang jumlahya ga sesuai
8	Username_9	gpt enak dipake asal tahu kalimat perintah yang benar
9	Username_10	palingan revisi lagi

Gambar 4.5 Tampilan Cleaning

Pada Gambar 4.5 ditampilkan hasil proses pembersihan, yang menghilangkan komentar yang mengandung karakter non-tekstual, seperti simbol, angka, tanda baca, emotikon, dan elemen lain yang tidak relevan. Data tersebut kemudian dinormalisasikan.

c. Tampilan Normalisasi

Tampilan normalisasi menampilkan hasil dari proses konversi kata-kata tidak baku menjadi kata baku sesuai dengan kaidah Bahasa Indonesia. Hasil dari tahap ini dapat dilihat pada Gambar 4.6.



Preprocessing: Normalisasi

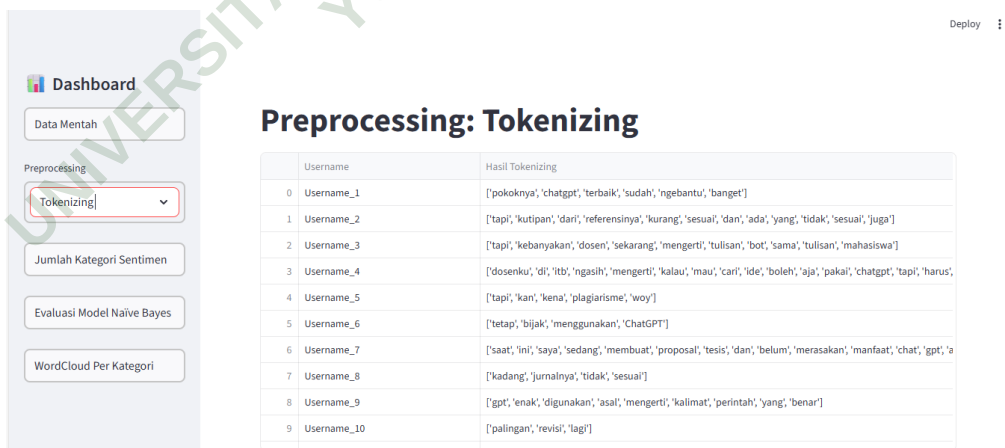
	Username	Hasil Normalisasi
0	Username_1	pokoknya chatgpt terbaik sudah ngebanget banget
1	Username_2	tapi kutipan dari referensinya kurang sesuai dan ada yang tidak sesuai juga
2	Username_3	tapi kebanyakan dosen sekarang mengerti tulisan bot sama tulisan mahasiswa
3	Username_4	dosenku di itb ngasih mengerti kalau mau cari ide boleh aja pakai chatgpt tapi harus di cek berkali dan sumi
4	Username_5	tapi kan kena plagiarisme woy
5	Username_6	tetap bijak menggunakan ChatGPT
6	Username_7	saat ini saya sedang membuat proposal tesis dan belum merasakan manfaat chat gpt apalagi cabang ilmu k
7	Username_8	kadang jurnalnya tidak sesuai
8	Username_9	gpt enak digunakan asal mengerti kalimat perintah yang benar
9	Username_10	palingan revisi lagi

Gambar 4.6 Tampilan Normalisasi

Pada Gambar 4.6 ditampilkan komentar-komentar yang telah dinormalisasi. Proses ini membuat data teks lebih konsisten dan siap untuk masuk ke tahap selanjutnya, yaitu *tokenizing*.

d. Tampilan Tokenizing

Tampilan tokenisasi menunjukkan hasil pemisahan kalimat menjadi kata-kata secara individual. Tujuan proses ini adalah mengonversi komentar menjadi unut kata yang lebih mudah diurai. Tampilan hasil *tokenizing* dapat di lihat pada Gambar 4.7.



Preprocessing: Tokenizing

	Username	Hasil Tokenizing
0	Username_1	['pokoknya', 'chatgpt', 'terbaik', 'sudah', 'ngebanget', 'banget']
1	Username_2	['tapi', 'kutipan', 'dari', 'referensinya', 'kurang', 'sesuai', 'dan', 'ada', 'yang', 'tidak', 'sesuai', 'juga']
2	Username_3	['tapi', 'kebanyakan', 'dosen', 'sekarang', 'mengerti', 'tulisan', 'bot', 'sama', 'tulisan', 'mahasiswa']
3	Username_4	['dosenku', 'di', 'itb', 'ngasih', 'mengerti', 'kalau', 'mau', 'cari', 'ide', 'boleh', 'aja', 'pakai', 'chatgpt', 'tapi', 'harus']
4	Username_5	['tapi', 'kan', 'kena', 'plagiarisme', 'woy']
5	Username_6	['tetapi', 'bijak', 'menggunakan', 'ChatGPT']
6	Username_7	['saat', 'ini', 'saya', 'sedang', 'membuat', 'proposal', 'tesis', 'dan', 'belum', 'merasakan', 'manfaat', 'chat', 'gpt', 'a']
7	Username_8	['kadang', 'jurnalnya', 'tidak', 'sesuai']
8	Username_9	['gpt', 'enak', 'digunakan', 'asal', 'mengerti', 'kalimat', 'perintah', 'yang', 'benar']
9	Username_10	['palingan', 'revisi', 'lagi']

Gambar 4.7 Tampilan Tokenizing

Pada Gambar 4.7 ditampilkan hasil proses *tokenizing* berupa daftar kata hasil pemisahan dari kalimat asli dalam komentar pengguna. Proses ini menjadi dasar penting sebelum dilakukan proses *stopword removal*.

e. Tampilan Stopword

Tampilan *stopword* digunakan menunjukkan hasil dari proses penghapusan kata-kata umum yang tidak memiliki pengaruh besar terhadap analisis sentimen. Tampilan hasil *stopword* dapat dilihat pada Gambar 4.8.



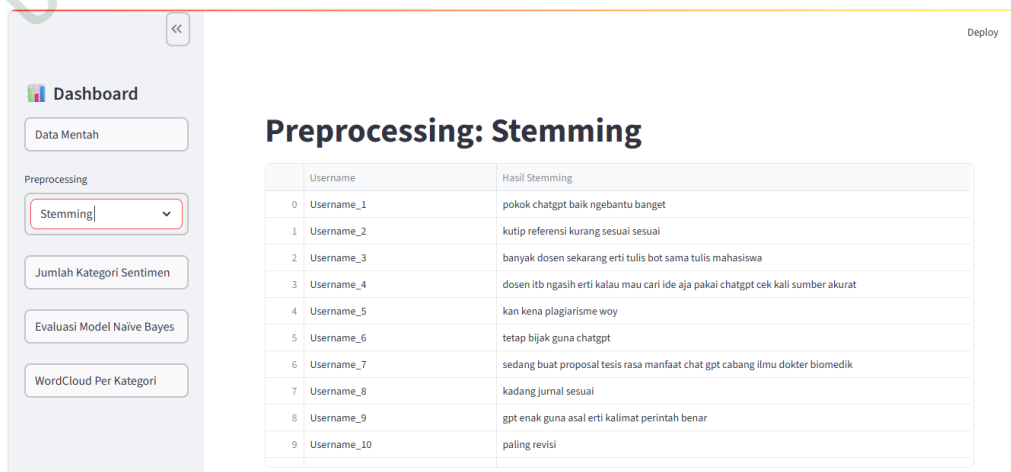
	Username	Hasil Stopword
0	Username_1	pokoknya chatgpt terbaik ngebantu banget
1	Username_2	kutipan referensinya kurang sesuai sesuai
2	Username_3	kebanyakan dosen sekarang mengerti tulisan bot sama tulisan mahasiswa
3	Username_4	dosenku itb ngasih mengerti kalau mau cari ide aja pakai chatgpt cek berkali sumbernya akurat
4	Username_5	kan kena plagiarisme woy
5	Username_6	tetap bijak menggunakan ChatGPT
6	Username_7	sedang membuat proposal tesis merasakan manfaat chat gpt cabang ilmu kedokteran biomedik
7	Username_8	kadang jurnalnya sesuai
8	Username_9	gpt enak digunakan asal mengerti kalimat perintah benar
9	Username_10	paling revisi

Gambar 4.8 Tampilan Stopword

Pada Gambar 4.8 ditampilkan visualisasi data komentar setelah proses *stopword removal*. Melalui tampilan ini, pengguna dapat melihat bagaimana teks menjadi lebih ringkas dan fokus pada kata-kata utama yang relevan, yang kemudian akan digunakan dalam proses analisis lanjutan seperti *stemming* atau klasifikasi sentimen.

f. Tampilan Stemming

Tampilan *stemming* menampilkan hasil data yang telah melalui proses *stemming*, yaitu proses mengubah kata berimbuhan menjadi bentuk dasar. Hasil *stemming* dapat dilihat pada Gambar 4.9.



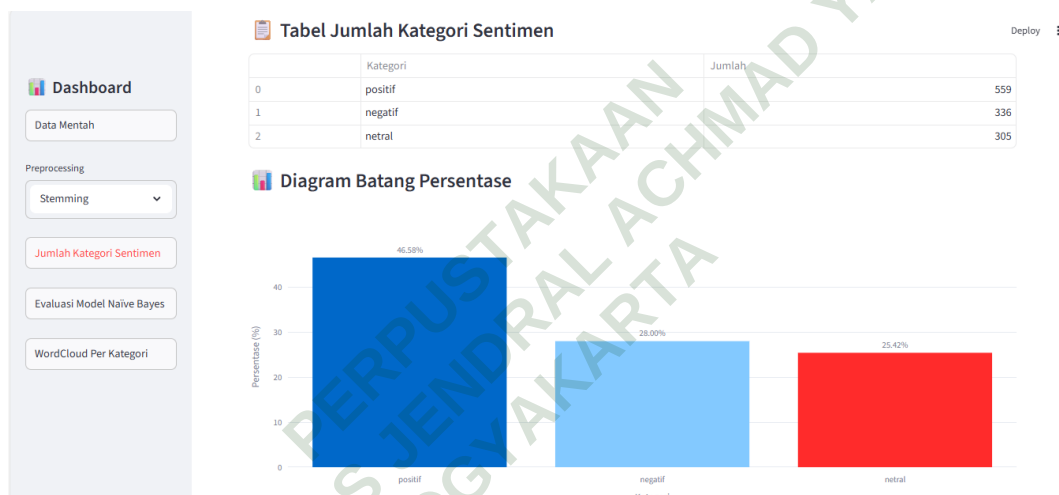
	Username	Hasil Stemming
0	Username_1	pokok chatgpt baik ngebantu banget
1	Username_2	kutip referensi kurang sesuai sesuai
2	Username_3	banyak dosen sekarang erti tulis bot sama tulis mahasiswa
3	Username_4	dosen itb ngasih erti kalau mau cari ide aja pakai chatgpt cek kali sumber akurat
4	Username_5	kan kena plagiarisme woy
5	Username_6	tetap bijak guna chatgpt
6	Username_7	sedang buat proposal tesis rasa manfaat chat gpt cabang ilmu dokter biomedik
7	Username_8	kadang jurnal sesuai
8	Username_9	gpt enak guna asal erti kalimat perintah benar
9	Username_10	paling revisi

Gambar 4.9 Tampilan Stemming

Pada Gambar 4.9 ditampilkan hasil stemming dari komentar pengguna TikTok yang telah diproses.

4.4.3 Tampilan Jumlah Kategori Sentimen

Tampilan jumlah kategori sentimen pada dashboard digunakan untuk menyajikan hasil klasifikasi sentimen terhadap komentar pengguna dalam bentuk tabel jumlah masing-masing kategori (positif, negatif, dan netral) serta diagram persentase yang dapat dilihat pada Gambar 4.10.



Gambar 4.10 Tampilan Jumlah Kategori Sentimen

Pada Gambar 4.10 ditampilkan hasil jumlah kategori sentimen dalam bentuk tabel dan diagram presentase. Kode berikut digunakan untuk menampilkan informasi tersebut.

```
df_label = pd.read_excel("labeling.xlsx")
if "Kategori" in df_label.columns:
    jumlah = df_label["Kategori"].value_counts()
    persentase = (jumlah / jumlah.sum() * 100).round(2)
    df_tabel = pd.DataFrame({
        "Kategori": jumlah.index,
        "Jumlah": jumlah.values
    })
    df_chart = pd.DataFrame({
        "Kategori": jumlah.index,
```

```

    "Persentase (%)": persentase.values
  })

  st.subheader("📊 Tabel Jumlah Kategori Sentimen")
  st.table(df_tabel)

  st.subheader("📊 Diagram Batang Persentase")
  fig = px.bar(df_chart, x="Kategori", y="Persentase (%)",
               text="Persentase (%)", color="Kategori")
  fig.update_traces(texttemplate='%{text:.2f}%',
                    textposition='outside')
  fig.update_layout(showlegend=False)
  st.plotly_chart(fig, use_container_width=True)

```

Kode di atas digunakan untuk menampilkan jumlah komentar berdasarkan kategori sentimen (positif, negatif, dan netral) dalam bentuk tabel dan diagram batang presentase pada dashboard.

4.4.4 Tampilan Evaluasi Model Naïve Bayes

Tampilan evaluasi model digunakan untuk menampilkan laporan hasil klasifikasi, dan hasil evaluasi performa dari algoritma Naïve Bayes yang digunakan dalam proses klasifikasi sentimen. Adapun kode program yang digunakan untuk menghasilkan tampilan laporan hasil klasifikasi adalah sebagai berikut.

```

st.title("📊 Evaluasi Model Naïve Bayes")
st.subheader(f"📊 Akurasi Model: {acc:.2%}")
st.subheader("📊 Tabel Laporan Klasifikasi")
df_report=pd.DataFrame(report).transpose().round(2)[["precision",
"recall", "f1-score"]]
st.dataframe(df_report, use_container_width=True)

```

Kode di atas digunakan untuk menampilkan hasil evaluasi performa model Naïve Bayes. Adapun tampilan laporan hasil klasifikasi yang dihasilkan dari kode tersebut dapat dilihat pada Gambar 4.11.

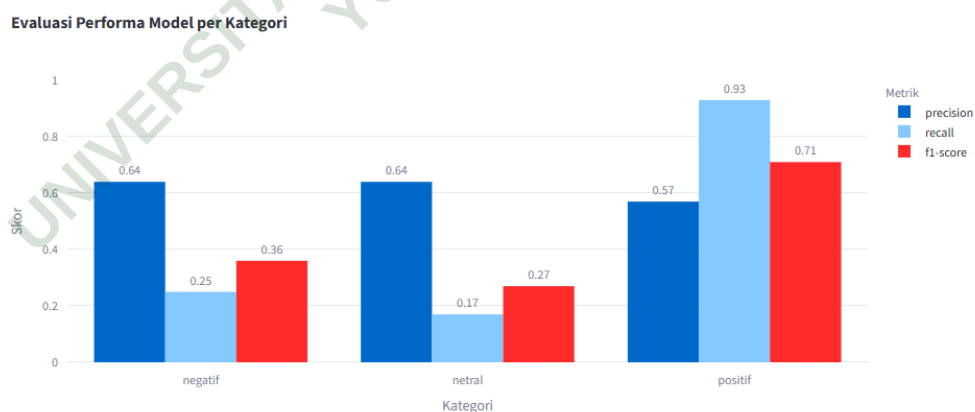
 Tabel Laporan Klasifikasi

	precision	recall	f1-score
negatif	0.64	0.25	0.36
netral	0.64	0.17	0.27
positif	0.57	0.93	0.71
accuracy	0.58	0.58	0.58
macro avg	0.62	0.45	0.45
weighted avg	0.61	0.58	0.52

Gambar 4.11 Tampilan Laporan Klasifikasi

Pada Gambar 4.11 menampilkan laporan hasil klasifikasi digunakan untuk menilai kinerja model Naïve Bayes dalam mengelompokkan sentimen dari komentar. Evaluasi ini mencakup tiga metrik utama yaitu *precision*, *recall*, dan *f1-score* untuk setiap kategori sentimen. Selain itu, disajikan pula nilai akurasi keseluruhan (*accuracy*) dari model, serta nilai *macro average* dan *weighted average* sebagai indikator performa rata-rata model terhadap seluruh kelas. Hasil evaluasi tersebut divisualisasikan dalam bentuk grafik yang ditampilkan pada Gambar 4.12.

 Diagram Evaluasi (Precision, Recall, F1-score)



Gambar 4.12 Tampilan Diagram Evaluasi (*Precision, Recall, F1-Score*)

Pada Gambar 4.12 menampilkan visualisasi evaluasi performa model klasifikasi dalam bentuk diagram batang. Diagram ini menggambarkan nilai *precision*, *recall*, dan *f1-score* untuk masing-masing kategori sentimen. Adapun

kode program yang digunakan untuk menghasilkan tampilan diagram evaluasi adalah sebagai berikut.

```
st.subheader("📊 Diagram Evaluasi (Precision, Recall, F1-score)")

df_plot = df_report[df_report.index.isin(y.unique())].reset_index().rename(columns={"index": "Kategori"})

fig_report=px.bar(df_plot.melt(id_vars="Kategori",value_vars=["precision", "recall", "f1-score"]),
                  x="Kategori",
                  y="value",
                  color="variable",
                  barmode="group",
                  labels={"value": "Skor", "variable": "Metrik"},
                  title="Evaluasi Performa Model per Kategori")

fig_report.update_layout(yaxis=dict(range=[0,1]),xaxis_title="Kategori", yaxis_title="Skor", legend_title="Metrik")

fig_report.update_traces(texttemplate='%{y}',textposition='outside')

st.plotly_chart(fig_report, use_container_width=True)

st.stop()
```

Kode di atas digunakan untuk menampilkan diagram batang yang memvisualisasikan hasil evaluasi performa model Naïve Bayes berdasarkan tiga metrik utama, *precision*, *recall*, dan *f1-score*.

4.5 PEMBAHASAN

Penelitian ini bertujuan untuk menganalisis sentimen komentar pengguna TikTok terhadap penggunaan ChatGPT dalam skripsi, dengan menerapkan algoritma Naïve Bayes untuk proses klasifikasi. Secara umum, hasil klasifikasi menunjukkan bahwa sebagian besar pengguna TikTok menunjukkan tanggapan yang positif terhadap penggunaan ChatGPT, disusul oleh tanggapan negatif dan netral. Hal ini mencerminkan bagaimana masyarakat digital, khususnya pengguna TikTok, merespon kemajuan teknologi kecerdasan buatan dalam dunia akademik.

Berdasarkan data sebanyak 1.200 komentar yang diperoleh melalui scraping menggunakan Tapicker, didapatkan proporsi sentimen sebagai berikut yang dapat dilihat pada tabel 4.5.

Tabel 4.5 Proporsi Sentimen

Positif	Negatif	Netral
47%	28%	25%

Tabel 4.5 menunjukkan bahwa sebagian besar komentar pengguna TikTok diklasifikasikan ke dalam kategori sentimen positif, yaitu sebesar 47% dari total data komentar yang dianalisis. Komentar positif umumnya mengandung apresiasi terhadap kemudahan yang diberikan oleh ChatGPT, seperti membantu menyusun kerangka penulisan, mencari ide judul, atau mempercepat proses pengerjaan skripsi. Hal ini mencerminkan bahwa sebagian besar pengguna memandang ChatGPT sebagai alat bantu yang bermanfaat dalam mendukung proses akademik, khususnya dalam penulisan tugas akhir.

Namun, sebanyak 28% komentar dikategorikan sebagai sentimen negatif, yang menunjukkan adanya kekhawatiran terhadap dampak penggunaan ChatGPT dalam konteks etika akademik. Beberapa komentar menyoroti risiko terjadinya plagiarisme, kualitas referensi yang tidak terverifikasi, serta kekhawatiran akan menurunnya keterampilan berpikir kritis mahasiswa karena terlalu bergantung pada teknologi. Sementara itu, 25% komentar dikategorikan sebagai netral. Sentimen ini umumnya tidak secara eksplisit menunjukkan dukungan maupun penolakan terhadap penggunaan ChatGPT. Komentar netral biasanya berupa pertanyaan, rasa penasaran, atau pernyataan informatif yang tidak memuat opini atau sikap tertentu.

Hasil klasifikasi menggunakan algoritma Naïve Bayes menunjukkan bahwa model memiliki akurasi sebesar 58,33%. Evaluasi performa model berdasarkan tiga metrik utama, yaitu *precision*, *recall*, dan *f1-score*, menghasilkan kinerja yang berbeda pada setiap kategori sentimen. Rincian evaluasi tersebut dapat dilihat pada Tabel 4.6.

Tabel 4.6 Evaluasi performa model

	Precision	Recall	F1-score
--	-----------	--------	----------

Negatif	0.64	0.25	0.36
Netral	0.64	0.17	0.27
Positif	0.57	0.93	0.71

Tabel 4.6 menunjukkan bahwa performa terbaik dicapai pada kategori sentimen positif, dengan *recall* sebesar 0.93 dan *f1-score* sebesar 0.71, menandakan bahwa model sangat baik dalam mengenali komentar-komentar positif. Sebaliknya, pada kategori negatif dan netral, meskipun *precision* masing-masing cukup tinggi 0.64, nilai *recall* yang rendah yaitu 0.25 dan 0.17 menunjukkan bahwa model kesulitan dalam mendeteksi komentar yang bersifat kritis atau netral. Hal ini dapat disebabkan oleh ketidakseimbangan distribusi data selama proses pelatihan model, sehingga model lebih terfokus pada pola-pola sentimen dominan, dalam hal ini sentimen positif. Meskipun akurasi yang diperoleh hanya sebesar 58,33%, algoritma Naïve Bayes tetap relevan dan layak digunakan dalam penelitian ini, terutama untuk studi eksploratif yang berfokus pada klasifikasi teks dalam jumlah besar.

Setelah proses klasifikasi dan evaluasi model, penelitian ini juga dilengkapi dengan *dashboard* visualisasi. *Dashboard* ini berfungsi sebagai media penyajian hasil data yang telah melalui proses pengolahan dan klasifikasi. Dibangun menggunakan Streamlit, *dashboard* ini dirancang untuk menampilkan hasil-hasil akhir dari tahapan penelitian dalam bentuk yang lebih terstruktur dan mudah dipahami. Tampilan yang disajikan mencakup data mentah sebelum dilakukan *preprocessing*, hasil dari setiap tahap *preprocessing*, serta hasil akhir klasifikasi berupa distribusi jumlah komentar pada masing-masing kategori sentimen. Selain itu, *dashboard* juga menampilkan visualisasi evaluasi performa model dalam bentuk tabel dan grafik, termasuk nilai *precision*, *recall*, *f1-score*, serta persentase distribusi sentimen.