

**ANALISIS SENTIMEN KOMENTAR SOSIAL MEDIA TERKAIT
PENERAPAN *E-PARKING* PONOROGO(PARKIR-GO)
DENGAN METODE BIDIRECTIONAL ENCODER
FROM TRANSFORMERS (BERT)**

SKRIPSI

Diajukan Sebagai Salah Satu Syarat
Untuk Memperoleh Gelar Sarjana Jenjang Strata Satu (S1)
Pada Program Studi Teknik Informatika Fakultas Teknik
Universitas Muhammadiyah Ponorogo



RIANAWATI
20533390

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS TEKNIK
UNIVERSITAS MUHAMMADIYAH PONOROGO
(2025)**

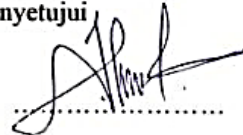
HALAMAN PENGESAHAN

Nama : Rianawati
NIM : 20533390
Program Studi : Teknik Informatika
Fakultas : Teknik
Judul Proposal Skripsi: Analisis Sentimen Komentar Sosial Media Terkait Penerapan E-Parking Ponorogo(PARKIR-GO) Dengan Metode Bidirectional Encoder From Transformers(BERT)

Isi dan formatnya telah disetujui dan dinyatakan memenuhi syarat
Untuk melengkapi persyaratan guna memperoleh Gelar Sarjana
pada Program Studi Teknik Informatika Fakultas Teknik Universitas Muhammadiyah
Ponorogo

Ponorogo, 07 Februari 2025

Menyetujui



Indah Puji Astuti, S.Kom., M.Kom

NIK. 19860424 201609 13

(Dosen Pembimbing Utama)

Ghulam Asrofi Buntoro, ST., M.Eng

NIK. 19870723 202109 12

Dosen Pembimbing Pendamping)



Mengetahui

Dekan Fakultas Teknik,

Ketua Program Studi Teknik Informatika,



(Edy Kurniawan, S.T., M.T)

NIK. 19771026 200810 12



(Adi Fajaryanto Cobantoro, S.Kom., M.Kom.)

NIK. 19840924 201309 13

PERNYATAAN ORISINILITAS SKRIPSI

Yang bertanda tangan di bawah ini:

Nama : Rianawati

NIM : 20533390

Program Studi: Teknik Informatika

Dengan ini menyatakan bahwa Skripsi saya dengan judul: “Analisis Sentimen Penerapan *E-Parking* Ponorogo(Parkir-Go) Dengan Metode Bidirectional *EncoderFrom* Transformer(BERT)” bahwa berdasarkan hasil penelusuran berbagai karya ilmiah, gagasan dan masalah ilmiah yang saya rancang/ teliti di dalam Naskah Skripsi ini adalah asli dari pemikiran saya. Tidak terdapat karya atau pendapat yang pernah ditulis atau diterbitkan oleh orang lain, kecuali yang secara tertulis dikutip dalam naskah ini dan disebutkan dalam sumber kutipan dan daftar pustaka.

Apabila ternyata di dalam Naskah Skripsi ini dapat dibuktikan terdapat unsur-unsur plagiatisme, saya bersedia Ijasah saya dibatalkan, serta diproses sesuai dengan peraturan perundang- undangan yang berlaku.

Demikian pernyataan ini dibuat dengan sesungguhnya dan dengan sebenar- benarnya.

Ponorogo, 24 Oktober 2024

Mahasiswa,



Rianawati

NIM. 20533390

HALAMAN BERITA ACARA UJIAN

Nama : Rianawati
NIM : 20533390
Program Studi : Teknik Informatika
Fakultas : Teknik
Judul Proposal Skripsi: Analisis Sentimen Komentar Sosial Media Terkait Penerapan E-Parking Ponorogo (PARKIR-GO) Dengan Metode Bidirectional Encoder From Transformers (BERT)

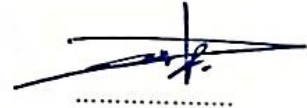
Telah diuji dan dipertahankan dihadapan
Dosen penguji tugas akhir jenjang Strata Satu(S1) pada:

Hari : Selasa
Tanggal : 31 Desember 2024
Dosen Penguji

Indah Puji Astuti, S.Kom., M.Kom
NIK. 19860424 201609 13
(Ketua Penguji)



Angga Prasetyo, S.T., M.Kom
NIK. 19820819 201112 13
(Anggota Penguji I)



Ir. Moh. Bhanu Setyawan, S.T., M.Kom
NIK. 19800225 201309 13
(Anggota Penguji II)



Mengetahui,

Dekan Fakultas Teknik,



(Edy Kurniawan, S.T., M.T)
NIK. 19771026 200810 12

Ketua Program Studi Teknik Informatika,



(Adi Fajaryanto Cobantoro, S.Kom., M.Kom.)
NIK. 19840924 201309 13

BERITA ACARA BIMBINGAN SKRIPSI

Nama : Rianawati
 NIM : 20533390
 Judul Skripsi : Implementasi Algoritma BERT dalam Sentimen
 Analisis Penerapan E-Parking
 Dosen Pembimbing I : Bu. Indah Puji Astuti, S.Kom., M.Kom

PROSES PEMBIMBINGAN

No	Tanggal	Materi Yang Dikonsultasikan	Saran Pembimbing / Hasil	Tanda Tangan
1	02/4 2024	Judul skripsi	- Ubah kasus yang diteliti	
2	05/4 2024	Judul skripsi	- Judul Disetujui - Lanjut BAB 1	
3	25/6 2024	BAB 1 2 3	- Penulisan nomor table, gambar, persamaan diperbaiki - Penambahan dasar teori terkait pengujian/evaluasi model BAB 2 - Rancangan GUI ditambahkan di BAB 3	
4	10/7 2024	Dataset BAB 1 2 3	- Pedoman Labelling Data ditambah - Melakukan percobaan modelling dataset ke bhs. Indo / Inggris	

No	Tanggal	Materi Yang Dikonsultasikan	Saran Pembimbing / Hasil	Tanda Tangan
5	11/7 2024	BAB 1 2	- Latar Belakang Ditingkas - Batasan masalah : - Sosial media - Total data - Poin 3 dihapus - Revisi tujuan - Penomoran rumus diperbaiki	
6	15/7 2024	BAB 3	+ Semua tabel dan gambar dirujuk sebelum ataupun sesudah - Visualisasi bagian modelling.	
7	16/7 2024	BAB 1 2 3	- Rumus BAB 3 dipindah di BAB 2	
8	17/7 2024	BAB 1 2 3	ACC Sempro	
9	1/10 2024	Bab 1-5	- tambah fitur untuk input data - tampilan desain web disesuaikan - tabel hasil percobaan - Saran di bab V	
10	2/10 2024	Aplikasi	Demo Aplikasi	

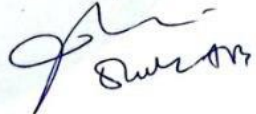
No	Tanggal	Materi Yang Dikonsultasikan	Saran Pembimbing / Hasil	Tanda Tangan
11	7/24 10	BAB 1-5	- Percobaan model diweb ditambah - Menjelaskan rincian tahap back translation. - Kesimpulan + saran diubah. - Uji coba lagi apakah aturan meningkat	
12	9/24 10	BAB 4-5	lanjut ke Gibitangan artikel	
13	18/24 10	+ Artikel - Abstrak - Naskah	- Kata Kunci/keyword urut abjad. pd abstrak - Ubah template jurnal ke jurnal yg sudah SINTA	
14	21/24 10	- Naskah - Skripsi	Abstrak Skripsi	
15	23/24 10	Artikel	Deep artikel	
16	24/24 10		Acc Sidang Skripsi	

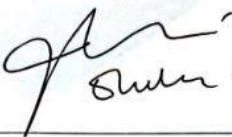
BERITA ACARA BIMBINGAN SKRIPSI

Nama : Rianawati
 NIM : 20533390
 Judul Skripsi : Implementasi Algoritma BERT dalam Sentimen Analisis Penerapan E-Parking Ponorego
 Dosen Pembimbing II : Bapak Ghulam Asrofi Buntoro, ST., M.Eng

PROSES PEMBIMBINGAN

No	Tanggal	Materi Yang Dikonsultasikan	Saran Pembimbing / Hasil	Tanda Tangan
1	01/7 2024	BAB 1 2 3	-BAB I disesuaikan dg format dlm Panduan. -Perbaiki penulisan sub-bab sentimen analisis -Perbaiki urutan Diagram/Gambar -Perbaiki Desain GUI	
2	10/7 24	BAB 1 BAB 2 BAB 3	-Perbaiki setiap baris baru dirapikan -Perbaiki penulisan Gambar.	
3	11/7 2024	BAB 1 2 3	- Indentasi Sub-bab diperbaiki - Spasi per paragraf belum rapi. (enter di setiap akhir Paragraf.)	
4	12/7 2024	BAB 1 2 3	- Masih ada tumor/persamaan yang belum dinomori.	

No.	Tanggal	Materi Yang Dikonsultasikan	Saran Pembimbing / Hasil	Tanda Tangan
5	13/7 2024	BAB 123	- Penambahan Kiri Gambar di BAB 3	
6	15/7 2024	BAB 123	- Ada paragraf yang tidak rapi di BAB 2	
7	16/7 2024	BAB 123	- Penambahan sitasi pada gambar dari sumber lain.	
8	17/07/2024		see sample  Dwison	
9	1/10 24	Aplikasi	Demo Aplikasi	
10	2/10 24	BAB 1-5	- Perbaiki penulisan Gambar di paragraf. - Penambahan fitur input teks	

No	Tanggal	Materi Yang Dikonsultasikan	Saran Pembimbing / Hasil	Tanda Tangan
11	7/10 2024	Bab 1-5	- Menambahkan penjelasan hasil uji coba dan tabel. -	
12	9/10 2024	Bab 4-5	- Merapikan kode di cobanya untuk persiapan - Lanjut mengerjakan Artikel	
13	14/10 2024	Artikel	- Penulisan disesuaikan lagi seperti bagian kesimpulan yg harusnya bentuk paragraf - Coba cari jurnal lain dan sesuaikan template (SMTA)	
14	18/10/2024		see below  Sulman	
15				
16				

SURAT KETERANGAN HASIL PLAGIASI SKRIPSI



UNIVERSITAS MUHAMMADIYAH PONOROGO LEMBAGA LAYANAN PERPUSTAKAAN

Jalan Budi Utomo No. 10 Ponorogo 63471 Jawa Timur Indonesia

Telp. (0352) 481124, Fax (0352) 461796, e-mail : lib@umpo.ac.id

website : www.library.umpo.ac.id

TERAKREDITASI A

(SK Nomor 000137/ LAP.PT/ III.2020)

NPP. 3502102D2014337

SURAT KETERANGAN HASIL *SIMILARITY CHECK* KARYA ILMIAH MAHASISWA UNIVERSITAS MUHAMMADIYAH PONOROGO

Dengan ini kami nyatakan bahwa karya ilmiah ilmiah dengan rincian sebagai berikut :

Nama : Rianawati

NIM : 20533390

Judul : ANALISIS SENTIMEN KOMENTAR SOSIAL MEDIA TERKAIT PENERAPAN E-PARKING PONOROGO(PARKIR-GO) DENGAN METODE BIDIRECTIONAL ENCODERFROM TRANSFORMER(BERT)

Fakultas / Prodi : Teknik Informatika

Dosen pembimbing :

1. Indah Puji Astuti, S.Kom., M.Kom
2. Ghulam Asrofi Buntoro, ST., M.Eng

Telah dilakukan check plagiasi berupa **SKRIPSI** di Lembaga Layanan Perpustakaan Universitas Muhammadiyah Ponorogo dengan prosentase kesamaan sebesar **17 %**

Demikian surat keterangan dibuat untuk digunakan sebagaimana mestinya.

Ponorogo, 14/02/2025

Kepala Lembaga Layanan Perpustakaan



Ayu Wulansari, S.Kom, M.A
NIK. 19760811 201111 21

SURAT KETERANGAN HASIL PLAGIASI ARTIKEL



**UNIVERSITAS MUHAMMADIYAH PONOROGO
LEMBAGA LAYANAN PERPUSTAKAAN**

Jalan Budi Utomo No. 10 Ponorogo 63471 Jawa Timur Indonesia
Telp. (0352) 481124, Fax (0352) 461796, e-mail : lib@umpo.ac.id

website : www.library.umpo.ac.id

TERAKREDITASI A

(SK Nomor 000137/ LAP.PT/ III.2020)

NPP. 3502102D2014337

SURAT KETERANGAN HASIL *SIMILARITY CHECK* KARYA ILMIAH MAHASISWA UNIVERSITAS MUHAMMADIYAH PONOROGO

Dengan ini kami nyatakan bahwa karya ilmiah ilmiah dengan rincian sebagai berikut :

Nama : Rianawati

NIM : 20533390

Judul : Analisis Sentimen terkait Penerapan E-Parking Ponorogo (Parkir-Go) dengan Metode Bidirectional Encoder from Transformers (BERT)

Fakultas / Prodi : Teknik Informatika

Dosen pembimbing :

1. Indah Puji Astuti, S.Kom., M.Kom
2. Ghulam Asrofi Buntoro, ST., M.Eng

Telah dilakukan check plagiasi berupa **Jurnal** di Lembaga Layanan Perpustakaan Universitas Muhammadiyah Ponorogo dengan prosentase kesamaan sebesar **11 %**

Demikian surat keterangan dibuat untuk digunakan sebagaimana mestinya.

Ponorogo, 14/02/2025

Kepala Lembaga Layanan Perpustakaan



Ayu Wulansari, S.Kom, M.A
NIK. 19760811 201111 21

MOTTO DAN PERSEMBAHAN

MOTTO

“Kunci bertahan dan bisa terus mengusahakan kehidupan lagi dan lagi adalah dengan menerima bahwa tidak ada sesuatu yang terjadi tanpa izin Allah SWT.

Everything happens for a reason”

“Just because it’s taking time doesn’t mean it’s not happening”

PERSEMBAHAN

Dengan mengucapkan syukur kepada Allah SWT, skripsi ini saya persembahkan sebagai tanda bukti kepada orang-orang yang telah senantiasa memberikan saya bantuan dan dukungan, keluarga saya bapak, ibu, nenek; kakak-kakak saya yang telah menjadi figure penting dalam proses tumbuh dan selalu membantu dalam berbagai bentuk baik psikis maupun materi; kepada sahabat-sahabat yang sering saya jadikan tempat berdiskusi dan tempat saya mencurahkan hal yang dialami; dan tentu saja paling penting adalah untuk diri sendiri yang masih mau mengusahakan segala sesuatu yang sudah dipilih dalam hidup walaupun banyak kurangnya.

ANALISIS SENTIMEN PENERAPAN *E-PARKING* PONOROGO(PARKIR-GO) DENGAN METODE BIDIRECTIONAL *ENCODER*FROM TRANSFORMER(BERT)

Rianawati, Indah Puji Astuti, Ghulam Asrofi Buntoro
Program Studi Teknik Informatika, Fakultas Teknik,
Universitas Muhammadiyah Ponorogo
e-mail: rianawati901@gmail.com

Abstrak

Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap kebijakan *E-Parking* di Kabupaten Ponorogo melalui media sosial serta mengukur akurasi kinerja algoritma BERT dalam analisis sentimen multibahasa. Hasil penelitian menunjukkan bahwa model BERT multilingual base memiliki performa yang kurang optimal, terutama dalam mengklasifikasikan sentimen positif pada data yang sebagian besar berbahasa jawa, dengan akurasi awal sebesar 50%. Setelah penerjemahan data ke Bahasa Indonesia dan Inggris, akurasi meningkat menjadi 55%, meskipun klasifikasi sentimen netral dan positif masih belum memuaskan. Penerapan metode augmentasi data melalui metode *back translation* menghasilkan peningkatan akurasi signifikan hingga 98% pada percobaan kedua. Analisis sentimen menunjukkan bahwa sentimen negatif terutama dipicu oleh kekhawatiran terkait kerumitan teknologi, khususnya bagi juru parkir yang sudah lanjut usia. Sentimen netral mencerminkan kebingungan mengenai metode pembayaran, sedangkan sentimen positif mendukung transparansi dan modernisasi yang dibawa oleh implementasi *E-Parking*, meski masih ada resistensi terhadap perubahan.

Kata Kunci: Analisis Sentimen, BERT, E-Parking, Media Sosial, Ponorogo

KATA PENGANTAR

Segala puji dan syukur penulis panjatkan ke hadirat Allah SWT atas limpahan rahmat, karunia, serta hidayah-Nya, sehingga penulis dapat menyelesaikan skripsi yang berjudul "*Analisis Sentimen Penerapan E-Parking Ponorogo (Parkir-Go) dengan Metode Bidirectional Encoder Representations from Transformers (BERT)*". Skripsi ini disusun sebagai salah satu syarat untuk memperoleh gelar Sarjana pada Program Studi Teknik Informatika, Fakultas Teknik, Universitas Muhammadiyah Ponorogo.

Penelitian ini bertujuan untuk menganalisis persepsi masyarakat terhadap penerapan sistem parkir elektronik *E-Parking* di Kabupaten Ponorogo melalui platform Parkir-Go. Metode *Bidirectional Encoder Representations from Transformers* (BERT) digunakan untuk mengklasifikasikan sentimen publik dari berbagai komentar di media sosial. Penulis berharap hasil penelitian ini dapat memberikan gambaran yang jelas mengenai respon masyarakat terhadap inovasi teknologi tersebut, serta memberikan kontribusi dalam evaluasi dan pengembangan layanan *E-Parking* di masa yang akan datang.

Dalam penyusunan skripsi ini, penulis menyadari bahwa terdapat banyak bantuan, bimbingan, dan dukungan yang telah diberikan oleh berbagai pihak. Oleh karena itu, pada kesempatan ini, penulis ingin menyampaikan rasa terima kasih yang mendalam kepada:

1. Yang terhormat bapak Edy Kurniawan S.T., M.T., selaku Dekan Fakultas Teknik Universitas Muhammadiyah Ponorogo.
2. Yang terhormat Bapak Adi Fajaryanto Cobantoro, S.Kom, M.Kom selaku Ka Prodi Teknik Informatika.
3. Ibu Indah Puji Astuti S.Kom, M.Kom selaku dosen pembimbing I, yang dengan sabar memberikan banyak bimbingan, arahan, saran, serta semangat dalam penyusunan skripsi ini hingga selesai.

4. Bapak Ghulam Asrofi Buntoro, ST., M.Eng selaku dosen pembimbing II, yang telah memberikan bimbingan dan arahan serta semangat kepada penulis selama proses penyusunan skripsi ini.
5. Keluarga tercinta yang senantiasa memberikan doa, serta dorongan positif yang tiada henti.
6. Rekan-rekan mahasiswa yang telah memberikan semangat dan kerjasama yang luar biasa selama masa studi terutama teman teman Angkatan 2020 dan kelas D yang memberikan pengalaman sebagai mahasiswa. Serta mungkin rekan-rekan yang saya kenal ketika berorganisasi terimakasih atas pengalaman dalam beberapa hal.
7. Sahabat – sahabat yang telah saya susahkan dan sering menjadi tempat berdiskusi, bercerita yang senantiasa memberikan dukungan.
8. Diri sendiri yang telah mengusahakan segala hal yang menjadi tanggungjawabnya. You did a great job!.

Penulis menyadari bahwa skripsi ini masih jauh dari kesempurnaan. Oleh karena itu, penulis sangat terbuka terhadap kritik dan saran yang membangun demi perbaikan dan pengembangan penelitian lebih lanjut. Semoga skripsi ini dapat bermanfaat bagi para pembaca dan pihak-pihak yang membutuhkan.

Ponorogo, 06 Oktober 2024

Rianawati

DAFTAR ISI

HALAMAN PENGESAHAN.....	i
PERNYATAAN ORISINILITAS SKRIPSI	ii
HALAMAN BERITA ACARA UJIAN	iii
BERITA ACARA BIMBINGAN SKRIPSI	iv
SURAT KETERANGAN HASIL PLAGIASI SKRIPSI	x
SURAT KETERANGAN HASIL PLAGIASI ARTIKEL	xi
MOTTO DAN PERSEMBAHAN	xii
Abstrak.....	xiii
KATA PENGANTAR.....	xiv
DAFTAR ISI	xvi
DAFTAR GAMBAR	xviii
DAFTAR TABEL	xxi
BAB I PENDAHULUAN.....	1
1.1. Latar Belakang	1
1.2. Rumusan Masalah	2
1.3. Batasan Masalah.....	3
1.4. Tujuan Penelitian.....	3
1.5. Manfaat Penelitian.....	3
BAB II TINJAUAN PUSTAKA	5
2.1 Penelitian Terdahulu.....	5
2.2 Landasan Teori	11
BAB III METODE PENELITIAN	34
3.1. Business Understanding	36

3.2.	Data Understanding.....	38
3.3.	Data Preparation.....	38
3.4.	Modeling	44
3.5.	Evaluation	48
3.6.	Deployment	49
BAB IV HASIL DAN PEMBAHASAN		52
4.1.	Data Collecting.....	52
4.2.	Pre-Processing Data	53
4.3.	Data Labelling.....	57
4.4.	Modelling	59
4.5.	Evaluasi Model.....	61
4.6.	Analisis dan Intepretasi	71
4.7.	Graphic User Interface	78
4.8.	Hasil Pengujian Model Klasifikasi.....	80
4.9.	Black Box Testing	86
BAB V PENUTUP.....		89
5.1.	Kesimpulan	89
5.2.	Saran	89
DAFTAR PUSTAKA.....		90

DAFTAR GAMBAR

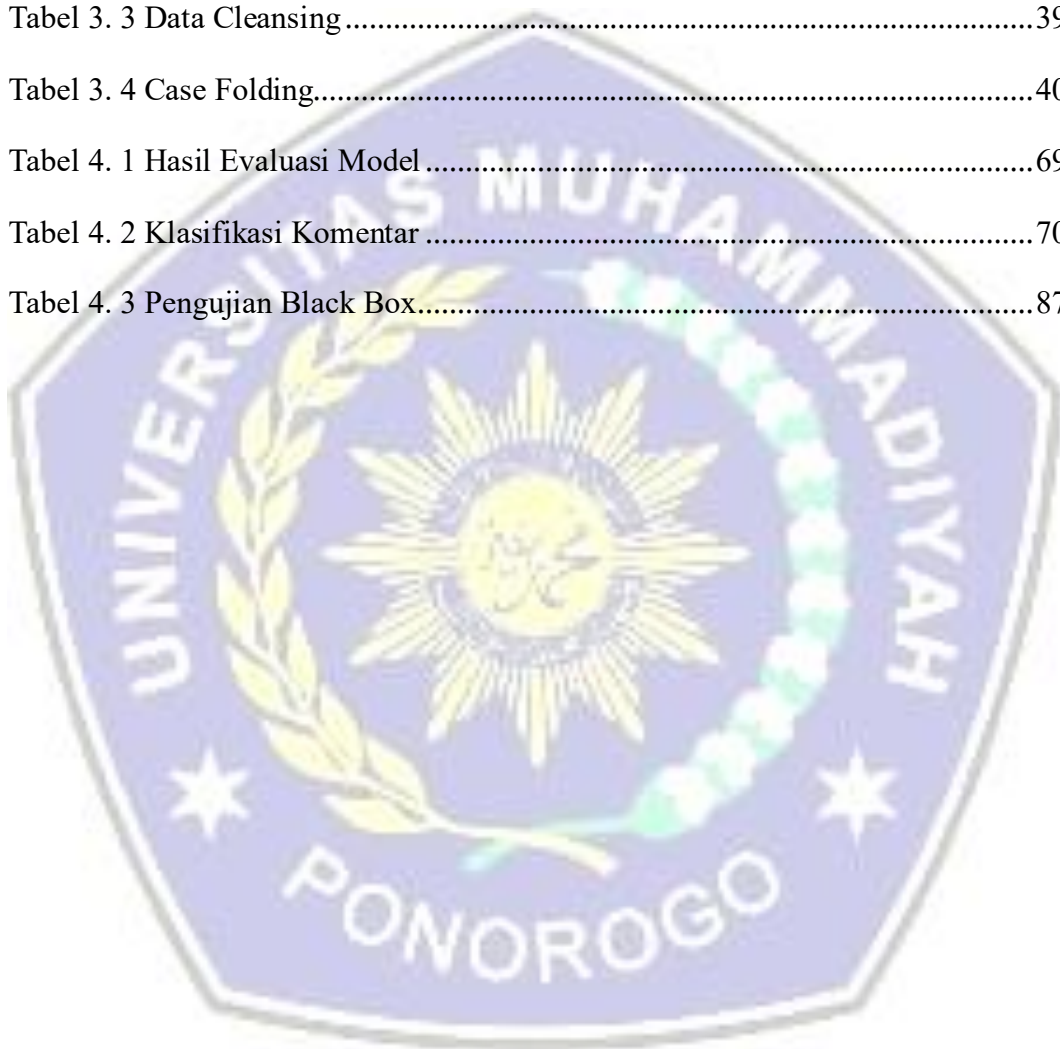
Gambar 2. 1 Blok Bangunan Bahasa dan Aplikasinya [19].....	14
Gambar 2. 2 Arsitektur Transformers[17].....	17
Gambar 2. 3 Contoh Mekanisme Self-Attention[18].....	18
Gambar 2. 4 Blok <i>Encoder</i> [17].....	19
Gambar 2. 5 Blok <i>Decoder</i> [17].....	22
Gambar 2. 6 Word Embending[20].....	25
Gambar 2. 7 <i>Pre-Training</i> dan Fine Tuning BERT [20].....	27
Gambar 3. 1 Diagram Alur Metodologi Penelitian	35
Gambar 3. 2 Dataset Hasil Scrapping.....	38
Gambar 3. 3 Token WordPiece	41
Gambar 3. 4 Token Embedding	41
Gambar 3. 5 Token Padding.....	42
Gambar 3. 6 Indeks ID token BERT	42
Gambar 3. 7 Tokenisasi ID BERT.....	43
Gambar 3. 8 Sentence Embedding.....	43
Gambar 3. 9 Positional Embedding	43
Gambar 3. 10 Proses Representasi <i>Input</i>	44
Gambar 3. 11 Tahap Klasifikasi.....	46
Gambar 3. 12 Ilustrasi Klasifikasi Sentimen BERT	48
Gambar 3. 13 Desain Halaman Dashboard.....	49
Gambar 3. 14 Desain Halaman Data Train	50

Gambar 3. 15 Desain Halaman Data Test	50
Gambar 3. 16 Desain Halaman Klasifikasi BERT	51
Gambar 3. 17 Halaman Visualisasi Data.....	51
Gambar 4. 1 Scrapping Komentar Instagram.....	52
Gambar 4. 2 Scrapping Komentar Youtube	52
Gambar 4. 3 Hasil Scrapping	53
Gambar 4. 4 Import Library Python.....	53
Gambar 4. 5 Missing Value Handling.....	54
Gambar 4. 6 Pembersihan Simbol.....	54
Gambar 4. 7 Hasil Cleansing	55
Gambar 4. 8 Stopword Removal.....	56
Gambar 4. 9 Slang Removal	56
Gambar 4. 10 Case Folding.....	57
Gambar 4. 11 Proses Labelisasi Data.....	58
Gambar 4. 12 Hasil Labelisasi Data.....	58
Gambar 4. 13 Inisiasi Model BERT pre-Training.....	59
Gambar 4. 14 Load Optimizer AdamW	60
Gambar 4. 15 Data Augmentasi	60
Gambar 4. 16 Evaluasi Model Dataset Multibahasa.....	62
Gambar 4. 17 Hasil Confusion matrix Dataset Multibahasa	62
Gambar 4. 18 Evaluasi Model Dataset Bahasa Inggris.....	63
Gambar 4. 19 Confusion matrix Dataset Bahasa Inggris.....	64
Gambar 4. 20 Evaluasi Model Dataset Bahasa Indonesia	65
Gambar 4. 21 Confusion matrix Dataset Bahasa Indonesia.....	65
Gambar 4. 22 Evaluasi Model Setelah Augmentasi Data(1)	67
Gambar 4. 23 Confusion matrix Data Augmentasi(1)	67
Gambar 4. 24 Evaluasi Model Setelah Augmentasi Data(2)	68
Gambar 4. 25 Confusion matrix Data Augmentasi(2)	68
Gambar 4. 26 Grafik modelling.....	70
Gambar 4. 27 Distribusi Sentimen.....	72

Gambar 4. 28 Distribusi Gender	72
Gambar 4. 29 Distribusi Sentimen berdasarkan Gender.....	73
Gambar 4. 30 WordCloud Sentimen Negatif	74
Gambar 4. 31 WordCloud Sentimen Netral	76
Gambar 4. 32 WordCloud Sentimen Positif.....	76
Gambar 4. 33 Halaman Dashboard.....	78
Gambar 4. 34 Halaman Data Train	79
Gambar 4. 35Halaman Data Test	79
Gambar 4. 36 Halaman Hasil Klasifikasi.....	80
Gambar 4. 37 Halaman Visualisasi Data.....	80
Gambar 4. 38 Teks Positif B.Indo.....	81
Gambar 4. 39 Teks Negatif B.Indo	81
Gambar 4. 40 Teks Positif B.Inggris.....	82
Gambar 4. 41 Teks Negatif B.Inggris	82
Gambar 4. 42 Teks Negatif B.Jawa.....	83
Gambar 4. 43 Teks Negatif B.Jawa.....	83
Gambar 4. 44 Teks Netral	84
Gambar 4. 45 Teks Netral 2	84

DAFTAR TABEL

Tabel 2. 1 Penelitian Terdahulu.....	5
Tabel 2. 2 Confusion Matrix	29
Tabel 3. 1 Waktu Penelitian.....	37
Tabel 3. 2 Data Collection.....	38
Tabel 3. 3 Data Cleansing	39
Tabel 3. 4 Case Folding.....	40
Tabel 4. 1 Hasil Evaluasi Model	69
Tabel 4. 2 Klasifikasi Komentar	70
Tabel 4. 3 Pengujian Black Box.....	87



BAB I

PENDAHULUAN

1.1. Latar Belakang

Pada tahun 2023, Pemerintah Kabupaten Ponorogo mengidentifikasi adanya kebocoran dalam setoran retribusi parkir. Untuk mengatasi permasalahan ini, Dinas Perhubungan dan BPPKD bekerja sama dengan Universitas Muhammadiyah Ponorogo dalam menerapkan uji coba sistem parkir elektronik PARKIR-GO. Sistem ini menggunakan aplikasi PARKIR-GO yang terpasang pada perangkat Electronic Data Capture (EDC) untuk menentukan jenis kendaraan dan tarif parkir sesuai ketentuan Peraturan Bupati Ponorogo No. 27 Tahun 2022. Uji coba ini dilakukan di beberapa lokasi strategis, seperti Jalan Hos Cokroaminoto, Jalan Suromenggolo, Aloon-Aloon Ponorogo, Jalan Gajah Mada, dan Jalan Sultan Agung, pada 22 Mei hingga 4 Juni 2023 dengan durasi 3 jam setiap harinya.

Namun, dalam penelitian pada Agustus 2023 di Pasar Tonatan dan Pasar Legi melalui pengamatan langsung menunjukkan adanya kendala teknis, seperti gangguan jaringan Wi-Fi yang menyebabkan mesin gagal mencetak tiket parkir. Selain itu, hasil evaluasi efektivitas dalam penelitian berdasarkan capaian retribusi dari tahun 2018 hingga 2023 menunjukkan bahwa sistem ini belum efektif. Pada tahun 2018-2022, target retribusi parkir mengalami penurunan hingga hanya mencapai 71% pada tahun 2022. Saat uji coba PARKIR-GO di 2023 dilakukan, target retribusi justru turun menjadi 61%, yang berarti belum memenuhi tujuan utama program[1].

Selain evaluasi berbasis capaian pendapatan, respons masyarakat terhadap kebijakan ini juga menjadi sarana penting dalam menilai keberhasilan implementasi dan bahan kajian lanjut. Selama masa uji coba, berbagai komentar muncul di media sosial, seperti Facebook, Instagram, dan YouTube. Beberapa akun yang membagikan informasi terkait PARKIR-GO menerima banyak interaksi, di antaranya InfoPonorogo (200 komentar), Ponorogo Update (267 komentar dan 17 komentar di unggahan lainnya), Gemasurya Facebook (22 komentar), serta KompasTV YouTube (32 komentar). Banyaknya tanggapan ini menunjukkan

adanya perhatian masyarakat terhadap kebijakan parkir elektronik, sehingga analisis sentimen dapat digunakan untuk memahami opini publik.

Analisis sentimen berbasis media sosial memungkinkan identifikasi opini, emosi, serta sikap masyarakat terhadap sistem PARKIR-GO. Metode ini dapat memberikan wawasan bagi pemerintah dalam mengevaluasi kebijakan dan menyusun perbaikan yang lebih sesuai dengan kebutuhan masyarakat. Evaluasi berbasis opini publik juga membuka ruang untuk kritik serta saran yang konstruktif guna meningkatkan efektivitas sistem parkir elektronik di Ponorogo.

Studi oleh Nanang Husin membandingkan algoritma Random Forest, Naïve Bayes, dan BERT untuk klasifikasi artikel berita CNN (2011-2022), dengan BERT mencatat akurasi tertinggi sebesar 92%[2]. Penelitian lain oleh Rhini Fatmasari et al. menguji SVM dan BERT untuk analisis sentimen komentar Twitter terkait layanan kampus, dengan BERT mencapai akurasi 90% dan SVM 80%[3]. BERT, dikenal dengan kemampuannya dalam memahami konteks bahasa alami secara *bidirectional*, telah terbukti efektif dalam menangani variasi linguistik dalam bahasa multikultural dan multibahasa[4][5]. Model-model NLP tradisional dan sekuensial seperti RNN dan LSTM seringkali mengalami keterbatasan dalam menangkap konteks yang kompleks dan menghadapi masalah *vanishing gradient* [6]. Dengan demikian, implementasi BERT dalam analisis sentimen terkait *E-Parking* di Ponorogo pada tahun 2023 menggunakan data komentar sosial media, karena keunggulan arsitekturnya yang efisien dan kemampuan mendalam dalam memahami konteks.

1.2. Rumusan Masalah

Berikut rumusan masalah yang menjadi dasar penelitian ini, yaitu:

1. Bagaimana analisis sentimen masyarakat terhadap implementasi sistem parkir elektronik PARKIR-GO di Kabupaten Ponorogo berdasarkan komentar di media sosial?
2. Bagaimana akurasi model BERT dalam mengklasifikasikan sentimen komentar terkait PARKIR-GO?

1.3. Batasan Masalah

Beberapa masalah ditetapkan untuk membatasi cakupan penelitian ini sehingga lebih terfokus dan terarah. Berikut batasan yang ditetapkan oleh peneliti dalam penelitian ini:

- 1) Data yang diambil berasal dari komentar pada konten sosial media meliputi Instagram, facebook, dan youtube yang berfokus pada uji coba dan penerapan *E-Parking* di Kabupaten Ponorogo.
- 2) Data dikumpulkan dengan teknik *web scraping* dalam rentang waktu Januari-Mei 2023.
- 3) Total data komentar yang dikumpulkan sejumlah 538 data, yang kemudian menjadi 527 data setelah proses *pra-processing*.
- 4) Penelitian berfokus pada penggunaan model BERT *base multilingual case* untuk analisis sentimen.
- 5) Metode BERT yang diterapkan menggunakan *pre-trained* BERT dengan *Fine-Tuning*.
- 6) Pemodelan dilakukan menggunakan Google Colab dan Python Versi 3.
- 7) Sentimen diklasifikasikan menjadi tiga kategori: positif, netral, dan negatif.

1.4. Tujuan Penelitian

Tujuan penelitian ini, sebagai berikut:

1. Menganalisis sentimen masyarakat terhadap implementasi *E-Parking* di Kabupaten Ponorogo menggunakan data komentar dari media sosial.
2. Mengevaluasi performa model BERT dalam mengklasifikasikan sentimen komentar terhadap *E-Parking* untuk menentukan keakuratannya dalam memahami opini publik.

1.5. Manfaat Penelitian

Manfaat yang akan diperoleh dengan penelitian ini diuraikan sebagai berikut:

1.1.1 Bagi Peneliti

Peneliti akan mendapatkan manfaat yang signifikan melalui pengembangan kemampuan analisis serta pengetahuan yang lebih dalam domain spesifik terkait *E-Parking* di Kabupaten Ponorogo, memberikan kontribusi pada pengembangan teknologi NLP melalui

implementasi algoritma BERT dalam analisis sentimen. Selain itu, peneliti akan memperoleh pengalaman berharga dalam merancang dan melaksanakan proyek penelitian yang dapat membantu dalam pengembangan karir mereka di masa depan.

1.1.2 Bagi Pengembang Sistem atau Pemerintah Kabupaten Ponorogo

Penelitian ini dapat memberikan beberapa manfaat bagi Pemerintah Daerah Kabupaten Ponorogo sebagai pelaku yang memberlakukan kebijakan, antara lain:

a. Pemahaman Sentimen Masyarakat

Hasil analisis sentimen terhadap persepsi masyarakat di sosial media melalui komentar instagram dapat dimanfaatkan untuk mengevaluasi efektivitas program dan menyesuaikan kebijakan atau strategi berdasarkan umpan balik masyarakat.

b. Evaluasi Kinerja E-Parking

Dengan memahami sentimen masyarakat, Pemda dapat mengidentifikasi aspek yang perlu dievaluasi atau ditingkatkan dari performa program ini.

c. Perencanaan dan Pengembangan Berkelanjutan

Informasi yang diperoleh dari analisis sentimen dapat menjadi masukan berharga untuk perencanaan dan pengembangan program-program terkait *E-Parking* di masa depan.

BAB II

TINJAUAN PUSTAKA

2.1 Penelitian Terdahulu

Penelitian terdahulu merupakan dasar yang sangat penting dalam proses penelitian, karena membantu memahami perkembangan topik yang sedang diteliti dan mengidentifikasi kesenjangan atau peluang untuk kontribusi baru. Penelitian terdahulu juga bisa berfungsi sebagai sumber inspirasi yang nantinya membantu pelaksanaan penelitian. Berikut daftar penelitian terdahulu yang dijadikan bahan referensi dijelaskan pada Tabel 2.1.

Tabel 2. 1 Penelitian Terdahulu

No	Peneliti (Tahun)	Judul	Hasil
1.	Rhini Fatmasaria, Riska Kurnia Septianib, Tuahta Hasiolan Pinemb, Dedik Fabiyantob, Windu Gatab (2023)	Implementasi Algoritma BERT Pada Komentar Layanan Akademik dan Non Akademik Universitas Terbuka di Media Sosial[3]	Pada penelitian ini dilakukan analisa terkait komentar sosial media twitter dan tiktok sebanyak 658 data dengan hasil 281 sentimen positif dan 404 data <i>sentimen</i> negatif menggunakan beberapa teknik yang diuji meliputi SVM dengan 88.34% akurasi, Naïve Bayes – MultinomialNB akurasi 87.37%, Naïve Bayes – GausiyanNB akurasi 87.37%, K-NN 85.43%, Decision tree 83% dan logistic regression 83% kemudian setelah dilakukan metode machine learning dilakukan klasifikasi dengan metode BERT yang

			menghasilkan akurasi sebesar 90%.
2.	Ardiansyah, Adika Sri Widagdo, Krisna Nuresa Qodri, Fachruddin Edi Nugroho Saputro, Nisrina Akbar Rizky P (2023)	Analisis Sentimen Terhadap Pelayanan Kesehatan Berdasarkan Ulasan Google Maps Menggunakan BERT[7]	Penelitian ini menggunakan model BERT indo-base-p1 untuk analisis sentimen dengan dataset 4228 data. Proses klasifikasi dilakukan dengan metode Lexicon yang mengkategorikan tweet menjadi tiga kategori: Negatif, Netral, dan Positif. Labeling menggunakan transformers Hugging Face dengan model RoBERTa berbahasa Indonesia menunjukkan bahwa sentimen

			positif adalah yang tertinggi dengan 2460 data (58.2%). Model menunjukkan akurasi tinggi pada data validasi (85%) dan testing (86%), menandakan kemampuan model dalam mengenali pola dengan baik. Namun, macro average yang lebih rendah (75% pada validasi dan 73% pada testing) menunjukkan beberapa kelas tidak diprediksi dengan baik. Weighted average yang tinggi dan konsisten (85% pada validasi dan 86% pada testing) menunjukkan bahwa kelas-kelas besar diprediksi dengan sangat baik.
3.	Bayu Kurniawan, Ahmad Ari Aldino, Auliya Rahman Isnain (2022)	Sentimen Analisis Terhadap Kebijakan Penyelenggara Sistem Elektronik (PSE) Menggunakan Algoritma Bidirectional Encoder Representations From Transformers (BERT)[8]	BERT diuji dengan parameter <i>batch size</i> 16 dan epoch 5. Hasil menunjukkan bahwa pada satu pengujian, BERT mencapai akurasi 69%, tetapi pada dua pengujian lainnya hanya mencapai 55%. Saat menggunakan BERT, akurasi dipengaruhi oleh keseimbangan data. Meskipun dataset yang seimbang lebih kecil, ia menghasilkan akurasi lebih tinggi (62%) dibandingkan dengan

			dataset yang tidak seimbang. Ketidakseimbangan data dapat mengurangi akurasi model karena model cenderung belajar lebih dari kelas mayoritas dan mengabaikan kelas minoritas.
4.	Muhammad Mahrus Zain, Rizky Nathamael Simbolon, Harlem Sulung dan Zaidan Anwar (2021)	Analisis Sentimen Pendapat Masyarakat Mengenai Vaksin Covid-19 Pada Media Sosial Twitter Dengan Robustly Optimized BERT Pretraining Approach[9]	Pada analisis sentimen, dilakukan pelabelan menggunakan model Indonesian RoBERTa Base <i>Sentimen Classifier</i> yang telah dilatih sebelumnya. Model Indonesian RoBERTa adalah varian dari RoBERTa yang telah melalui pelatihan dengan memanfaatkan 527MB teks dari Wikipedia berbahasa Indonesia menggunakan teknik Masked Language Modeling (MLM). Dari total 109.202 data yang digunakan, model Indonesian RoBERTa Base <i>Sentiment Classifier</i> memprediksi 6.250 sentimen positif, 75.883 sentimen netral, dan 26.934 sentimen negatif. Akurasi keseluruhan dari prediksi yang dilakukan adalah sebesar 95%. Secara rata-rata, hasil akurasi prediksi untuk masing-masing label adalah sebagai berikut: 84% untuk

			sentimen positif, 97% untuk sentimen netral, dan 93% untuk sentimen negatif.
5.	Alwi Jaya (2023)	Analisis Sentimen Pandangan Publik Terhadap Profesi PNS (Pegawai Negeri Sipil) Dari Twiter Menerapkan Indonesian Roberta Base Sentimen Classifier[10]	Dari 394 data tweet yang telah diberi label, 4.8% dianggap positif, 8.6% dianggap negatif, dan sisanya, yaitu 86.4%, dianggap netral. Selain itu, setelah dilakukan uji coba pelabelan pada penelitian ini, ternyata prediksi Indonesian Roberta Base Sentiment Classifier yang dilakukan memiliki akurasi yang baik, dengan rata-rata keseluruhan akurasi sebesar 90%.
6.	Nanang Husin (2023)	Komparasi Algoritma Random Forest, Naïve Bayes, Dan Bert Untuk Multi-Class Classification Pada Artikel Cable News Network (CNN)[2]	Model BERT menunjukkan kinerja terbaik dalam mengklasifikasikan dataset artikel berita CNN dari tahun 2011 hingga 2022, yang terdiri dari 37.904 baris artikel. Dataset ini tidak seimbang karena jumlah artikel dalam setiap kategori memiliki perbedaan yang signifikan. Untuk menyeimbangkan data, penelitian ini menggunakan library "imblearn" dengan metode RandomUnderSampler.

			Algoritma BERT mencapai akurasi <i>training</i> sebesar 93% dan akurasi <i>testing</i> sebesar 92%, dengan marco avg dari f1 score sebesar 92%. Dengan demikian, algoritma BERT terbukti efektif dalam mengklasifikasikan teks artikel berita, terutama dalam kasus dataset yang cukup besar dan tidak seimbang dibanding metode random forest yang memiliki akurasi <i>training</i> 81% dan Naïve Baiyes 78%.
7.	Yono Cahyono, Saprudin (2019)	Analisis <i>Sentimen</i> Tweets Berbahasa Sunda Menggunakan Naive Bayes Classifier dengan Seleksi Feature Chi Squared Statistic[4]	Dataset yang digunakan terdiri dari 316 tweet yang mencampurkan bahasa Sunda dan bahasa Indonesia yang diperoleh dari Twitter melalui proses pengambilan data menggunakan RapidMiner. Tahap pre-processing dilakukan dengan melakukan case folding, tokenize, dan penghapusan stopword. Seleksi fitur menggunakan statistik chi square untuk memilih kata-kata yang penting untuk analisis dokumen, dengan melakukan optimasi seleksi menggunakan forward selection. Hasil penelitian

			menunjukkan bahwa penggunaan seleksi fitur Chi Square Statistic dan algoritma Naïve Bayes Classifier menghasilkan akurasi sebesar 78.48%.
--	--	--	---

2.2 Landasan Teori

2.2.1 Parkir Elektronik Ponorogo (PARKIR-GO)

Parkir Elektronik Ponorogo (PARKIR-GO) menerapkan konsep retribusi parkir tepi jalan umum (PARKIR-GO), yang merupakan sebuah sistem digitalisasi perparkiran yang berbasis android. Tujuannya adalah untuk memfasilitasi transaksi perparkiran dan menyederhanakan proses pendataan, sehingga transparansi pengeluaran masyarakat terhadap pendapatan daerah dari hasil parkir dapat termonitor secara *real-time*. Metode yang digunakan dalam pengembangan ParkirGo mengadopsi tahapan metode pengembangan sistem *waterfall*, yang meliputi analisis, perancangan, implementasi, dan pengujian. Hasilnya adalah aplikasi sistem parkir online berbasis android yang responsif, yang mampu beradaptasi dengan baik pada layar smartphone yang berukuran besar atau kecil. Aplikasi ini juga memastikan bahwa setiap pengguna mendapatkan struk parkir setelah melakukan parkir di tepi jalan[1].

Langkah ini dimaksudkan untuk meningkatkan manajemen parkir, menaikkan tarif parkir, dan pada akhirnya berperan dalam penataan fasilitas parkir di daerah tersebut. Dengan memanfaatkan teknologi, terutama melalui aplikasi Parkir-Go, pemerintah berusaha beralih dari sistem parkir manual ke sistem elektronik yang lebih efisien dan transparan. Penerapan parkir elektronik tidak hanya dimaksudkan untuk meningkatkan efisiensi layanan parkir, tetapi juga untuk mengatasi masalah seperti parkir ilegal dan pelanggaran kendaraan. Dengan menerapkan tarif parkir progresif, sesuai dengan Peraturan Bupati Kabupaten Ponorogo No.27 Tahun 2022, pemerintah berupaya

memastikan pungutan yang adil dan standar berdasarkan jenis kendaraan. Dinas Perhubungan memainkan peran penting dalam mengelola dan mengumpulkan biaya parkir, dengan tujuan meningkatkan pengelolaan pendapatan dan mendukung pembangunan ekonomi. Langkah penerapan sistem parkir elektronik didorong oleh berbagai faktor, termasuk kurangnya kedisiplinan masyarakat dan petugas parkir, rendahnya pendapatan daerah dari biaya parkir, dan prevalensi pungutan liar. Inisiatif pemerintah untuk menerapkan parkir elektronik awalnya difokuskan pada pasar-pasar utama seperti Pasar Tonatan dan Relokasi Pasar Legi. Langkah ini diharapkan tidak hanya akan meningkatkan pendapatan, tetapi juga akan meningkatkan efisiensi dan integritas keseluruhan proses pengumpulan biaya parkir.[1]

2.2.2 Sentimen Analisis

Sentimen analysis secara luas merujuk pada bidang komputasi linguistik, pengolahan bahasa alami, dan text mining dengan tujuan untuk menentukan sikap dari pembicara atau penulis terkait topik tertentu. Teknik ini digunakan untuk mengekstraksi informasi mengenai opini dan sentimen. Sentimen sendiri adalah perasaan seseorang terhadap suatu hal. Pengelompokan polaritas teks dalam kalimat, dokumen, atau fitur entitas untuk menentukan apakah pendapat yang disampaikan bersifat positif, negatif, atau netral adalah tugas utama dalam analisis sentimen.[5]

2.2.3 Text Mining

Text mining, juga dikenal sebagai teks analitik, merupakan proses mengubah teks tidak terstruktur menjadi data terstruktur untuk analisis yang lebih mudah. Ini melibatkan teknik-teknik seperti pemrosesan bahasa alami (NLP), pengenalan entitas bernama, analisis sentimen, dan ekstraksi informasi. Tujuan utama dari *text mining* adalah untuk memperoleh wawasan berharga dari teks yang tidak terstruktur, yang dapat digunakan

untuk berbagai aplikasi seperti analisis tren pasar, manajemen hubungan pelanggan, dan deteksi penipuan.

Teknik *text mining* sering kali dimulai dengan tahap pra-pemrosesan, yang meliputi pembersihan data, tokenisasi, dan normalisasi. Pembersihan data mencakup penghapusan karakter khusus, tanda baca, dan kata-kata yang tidak memiliki makna penting, seperti kata sambung. Tokenisasi adalah proses memecah teks menjadi unit-unit yang lebih kecil, seperti kata atau frasa. Normalisasi melibatkan pengubahan bentuk kata menjadi bentuk dasar atau lematinya.

Dalam *text mining* terdapat proses yang bertujuan mengubah data teks ke dalam bentuk numerik agar dapat diolah dan dianalisis di proses selanjutnya yakni *pre-processing* dengan tahapan sebagai berikut:

1) Case Folding

Teknik yang digunakan untuk mengubah seluruh huruf dalam teks menjadi huruf kecil atau *lowercase* untuk memudahkan pemrosesan teks.

2) Filtering

Proses ini melakukan penghapusan atau pengecualian dokumen atau kata-kata tertentu dari teks berdasar aturan dan kriteria, tujuannya memperbaiki kualitas relevansi data serta mempermudah proses analisis.

3) Stemming

Tahap Dimana kata diubah ke akar katanya atau kembali ke kata dasar.

4) Tokenizing

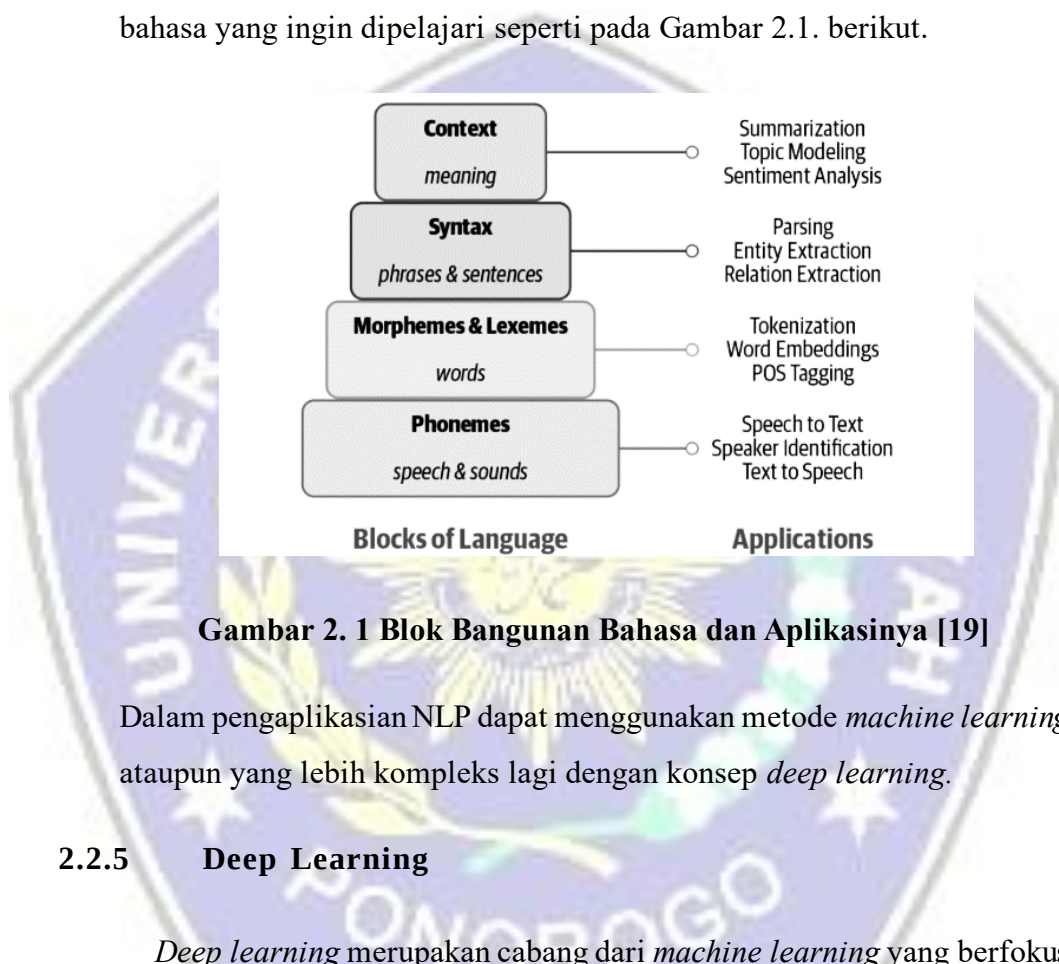
Proses ini membagi teks menjadi kata dan frasa, kalimat atau bagian bermakna.

2.2.4 Natural language Processing (NLP)

Natural language Processing atau NLP merupakan bidang yang berfokus pada pengembangan teknologi untuk memproses bahasa alami,

seperti bahasa inggris dan lainnya. Bidang ini adalah persimpangan dari *computer science*, *artificial intelligence* dan *linguistics*.

Dalam pengaplikasian NLP pada bahasa alami diperlukan pemahaman terkait struktur bahasa. Bahasa memiliki struktur sistem komunikasi yang kompleks dengan kombinasi komponen yang konstituen seperti karakter, kata, dan lainnya. Penerapan NLP pun bermacam-macam tergantung blok bahasa yang ingin dipelajari seperti pada Gambar 2.1. berikut.



Gambar 2. 1 Blok Bangunan Bahasa dan Aplikasinya [19]

Dalam pengaplikasian NLP dapat menggunakan metode *machine learning* ataupun yang lebih kompleks lagi dengan konsep *deep learning*.

2.2.5 Deep Learning

Deep learning merupakan cabang dari *machine learning* yang berfokus pada algoritma yang terinspirasi oleh struktur dan fungsi otak manusia, yang dikenal sebagai jaringan saraf tiruan. Dalam dekade terakhir, *Deep learning* telah menjadi salah satu teknologi paling inovatif dalam bidang kecerdasan buatan, memungkinkan terobosan signifikan dalam berbagai aplikasi seperti pengenalan gambar, pemrosesan bahasa alami, dan permainan komputer. *Deep learning* menggunakan berbagai arsitektur jaringan saraf, seperti *Multilayer Perceptrons* (MLP), *Convolutional*

Neural Networks (CNN), dan *Recurrent Neural Networks* (RNN). MLP merupakan jaringan dasar dengan lapisan *input*, satu atau lebih lapisan tersembunyi, dan lapisan *output*. CNN sangat efektif untuk pengenalan gambar dan video, menggunakan lapisan konvolusi untuk mendeteksi fitur visual, sementara RNN cocok untuk data deret waktu dan pemrosesan bahasa alami, menggunakan lapisan yang dapat mempertahankan informasi urutan[11].

Algoritma *Deep learning* menggunakan berbagai metode pembelajaran, termasuk pembelajaran terbimbing (*supervised learning*), pembelajaran tak terbimbing (*unsupervised learning*), dan pembelajaran penguatan (*reinforcement learning*). Dalam pembelajaran terbimbing, jaringan dilatih dengan dataset berlabel untuk memprediksi *output* yang benar. Pembelajaran tak terbimbing mencoba menemukan pola dalam data tanpa label, sedangkan pembelajaran penguatan mengharuskan jaringan belajar dengan mencoba-coba dan menerima umpan balik berupa *reward* atau *punishment*. [12]

Konsep dasar deep learning:

1. *Neural Network*

Konsep ini terdiri dari unit-unit pemrosesan(neuron) yang diorganisir dalam lapisan(layers). *Input* diberikan pada layer *input* dan *output* dihasilkan pada layer *output* setelah melalui serangkaian layer tersembunyi yang memungkinkan model untuk memahami fitur yang semakin kompleks pada data[6].

2. *Deep Neural Network*

Deep neural network mempunyai lebih dari satu layer tersembunyi. Semakin banyak layer menandakan kompleksnya model dan kemampuan yang baik dalam representasi yang abstrak dari data.

3. *Training*

Model dilatih dengan dataset besar untuk belajar menyesuaikan parameter agar memperoleh prediksi yang akurat.

4. *Backpropagation*

Algoritma ini menyesuaikan bobot(weight) dalam neural *network* berdasar perbedaan antara prediksi model dan nilai sebenarnya dari data *training*[12].

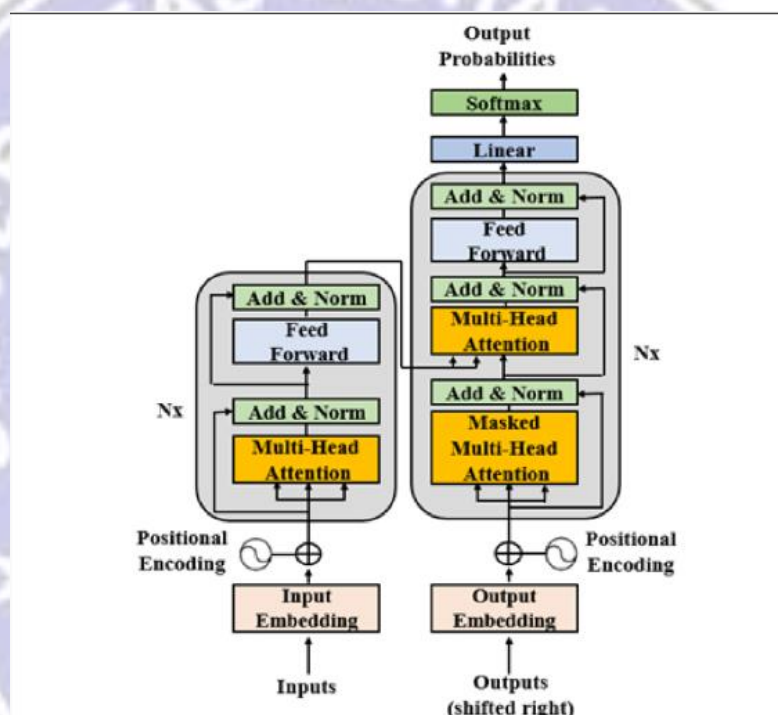
Deep learning kini telah banyak diterapkan dalam berbagai tugas pemrosesan salah satunya pemrosesan bahasa atau NLP. Kemampuan metode *Deep learning* untuk mempelajari berbagai masalah dengan konsep dan struktur mengikuti fungsi otak manusia memberikan kemampuan penyelesaian tugas yang baik. Bahasa adalah data yang kompleks dan tidak terstruktur. Pada tugas NLP membutuhkan model dengan representasi dan kemampuan belajar yang lebih baik untuk memahami dan menyelesaikan masalah bahasa[13].

2.2.6 Transformers

Vaswani,dkk pada 2017 memperkenalkan model *transformers* sebagai salah satu metode *Deep learning* dalam tugas pemrosesan bahasa alami atau NLP[14]. Ide inti dari *transformer* adalah mekanisme perhatian, yang memungkinkan model membaca kalimat dan memberikan perhatian lebih pada kata-kata tertentu. Saat memproses sebuah kata, *Transformer* memberikan "perhatian" pada kata-kata lain dalam kalimat tersebut. Untuk melakukan ini, Transformer menggunakan embedding yang mengubah kata-kata atau token ke dalam ruang vektor numerik. Dengan mekanisme *multihead attention*, model ini dapat memperhatikan hubungan antara kata-kata secara paralel, sehingga meningkatkan efisiensi. Mekanisme *self-attention* memungkinkan model memberikan bobot perhatian yang dinamis pada setiap kata dalam urutan *input*, tergantung pada konteksnya.

Model transformater didasarkan sepenuhnya pada mekanisme perhatian dan sepenuhnya menghilangkan reccurance. Metode ini menggunakan jenis mekanisme perhatian khusus yang disebut perhatian diri atau *self-attention*. Model transduksi urutan saraf yang bersaing umumnya mengadopsi struktur *encoder-decoder*. Model transduksi urutan saraf adalah sistem yang dapat mengubah urutan data saraf menjadi bentuk lain yaitu alat atau program

komputer yang dapat mengubah informasi dari satu bentuk urutan saraf ke bentuk urutan lainnya. Dalam sistem ini, bagian "*encoder*" bertugas mengubah urutan simbol-simbol *input* menjadi bentuk representasi kontinu yang disebut *z*. Setelah mendapatkan representasi *z*, "*decoder*" kemudian menghasilkan urutan *output* simbol-simbol satu per satu. Model ini bekerja secara *auto-regressive*, artinya pada setiap langkahnya, model menggunakan simbol yang dihasilkan sebelumnya sebagai masukan tambahan saat menghasilkan simbol berikutnya. Arsitektur Transformer mengadopsi pendekatan ini, dengan menggunakan perhatian diri bertumpuk dan lapisan-lapisan penuh yang terhubung pada *encoder* dan *decoder*.



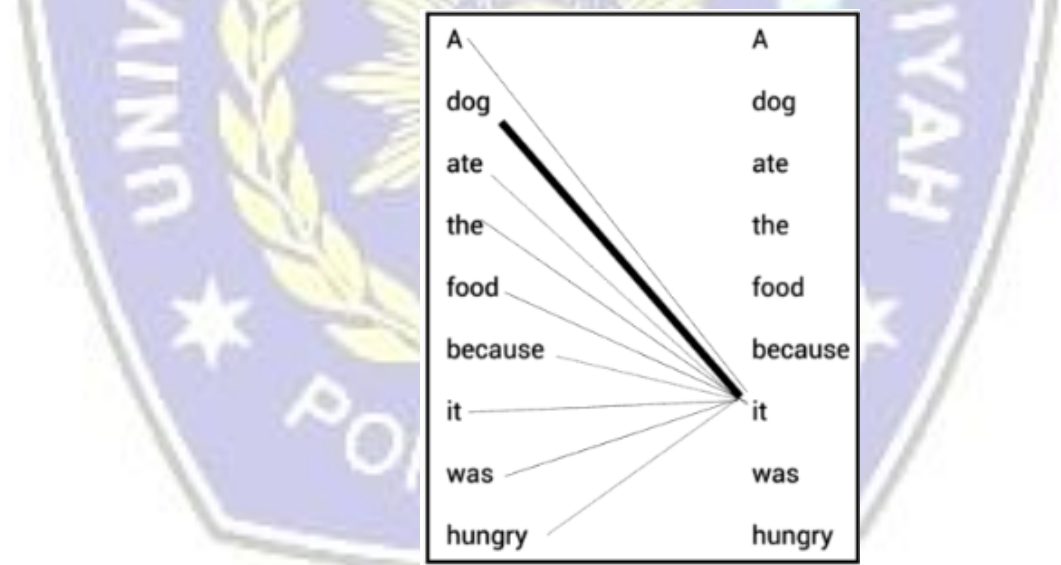
Gambar 2. 2 Arsitektur Transformers[14]

Berdasarkan Gambar 2.2 di sebelah kiri, *input* masuk ke sisi *encoder* Transformer melalui sublapisan perhatian dan sublapisan *FeedForward Network* (FFN). Di sebelah kanan, *output* target masuk ke sisi *decoder* Transformer melalui dua sublapisan perhatian dan sublapisan FFN. Dapat dilihat bahwa tidak ada RNN, LSTM, atau CNN.

Perhatian telah menggantikan pengulangan, yang memerlukan peningkatan jumlah operasi seiring bertambahnya jarak antara dua kata. Mekanisme perhatian adalah operasi "kata ke kata". Mekanisme perhatian akan menemukan bagaimana setiap kata terkait dengan semua kata lain dalam suatu urutan, termasuk kata yang sedang dianalisis itu sendiri.

- Mekanisme *Self-attention*

Manusia dapat dengan mudah menangkap makna dari setiap kata seperti kata ganti atau sebagainya mempresentasikan bagian yang mana. Namun model butuh mekanisme yang relevan agar dapat memahami maksud atau konteks dari kata. Dalam transformer mekanisme *self-attention* membantu model memahami kata melalui cara pengecekan relasional dari suatu kata ke setiap kata dalam kalimat[15]. Contoh kerja dari mekanisme ini seperti Gambar 2.3 Contoh Mekanisme *Self-attention* berikut.

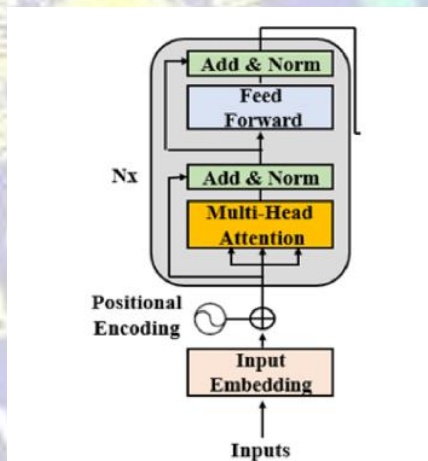


Gambar 2. 3 Contoh Mekanisme Self-Attention [18]

Mekanisme *self-attention* pada Gambar 2.3. melakukan perhatian terhadap kata “it” dengan mengecek ke seluruh kata dan menangkap maksud dari kata tersebut mengarah ke “dog”.

2.2.6.1 EncoderStack

Encoder pada Gambar 2.4 merupakan bagian penting dalam model *transformers* Gambar 2. 2. *Encoder* adalah komponen dalam dunia komputasi dan pemrograman yang berfungsi untuk menerjemahkan data dari suatu bentuk ke bentuk lain. *Transformer* terdiri dari tumpukan sejumlah N *encoder*. *Output* dari satu *encoder* dikirim sebagai *input* ke *encoder* di atasnya. Seperti yang ditunjukkan pada gambar berikut, terdapat setumpuk N jumlah *encoder*. Setiap *encoder* mengirimkan *output* nya ke *encoder* di atasnya. Format baru hasil *encoder* dalam *Deep learning* memungkinkan data diproses dan dikompresi oleh lapisan hilir. *Encoder* dalam *transformer* memiliki $N=6$ lapisan identik dengan dua sub-lapisan di setiap lapisan *encoder*. Pada sub-lapisan pertama adalah mekanisme perhatian diri *multi-head*, kemudian sub lapisan kedua merupakan jaringan *feed-forward* yang sederhana. *Encoder* akhir mengembalikan representasi kalimat sumber yang diberikan sebagai *output*.



Gambar 2. 4 Blok *Encoder* [17]

- *Input embedding*

Sublapisan penyematan berfungsi seperti model transduksi standar lainnya. Sebuah *tokenizer* akan mengubah kalimat menjadi token. Setiap *tokenizer* memiliki metodenya sendiri, tetapi hasilnya serupa. Misalnya, sebuah *tokenizer* yang

diterapkan pada urutan kalimat “Kebijakan yang positif dan mampu mensejahterakan.” Akan menghasilkan token berikut: ['kebijakan', 'yang', 'positif', 'dan', 'mampu', 'mensejahterakan'] Pada layer ini teks akan diubah dengan memisahkan masing-masing suku kata dan juga menjadi huruf kecil. Selanjutnya untuk proses embending token teks akan di ubah menjadi token id numerik. Token IDs = [101, 17710, 5638, 18317, 2319, 8675, 13433, 28032, 10128, 4907, 5003, 8737, 2226, 2273, 3366, 18878, 14621, 9126, 102]. Teks yang ditokenisasi harus disematkan(*embedding*). Banyak metode penyematan yang dapat diterapkan pada masukan yang ditokenisasi.

- Positional embedding

Tahapan selanjutnya adalah *positional embedding* untuk menggambarkan posisi token dalam suatu urutan. Tujuannya agar proses pemahaman teks sesuai urutan seharusnya.

- Sub-layer 1 multi-head attention

Sub-layer multi-head attention terdiri dari delapan *heads* dan diikuti oleh normalisasi *post-layer*, yang akan menambahkan koneksi residual ke *output* dari *sub-layer* dan menormalkannya. *Input* dari *sub-layer multi-attention* pada layer pertama dari tumpukan *encoder* adalah sebuah vektor yang berisi *embedding* dan *encoding* posisi dari setiap kata. Layer-layer berikutnya dari tumpukan ini tidak mengulangi operasi ini. Setiap kata dipetakan ke semua kata lain untuk menentukan bagaimana kata tersebut cocok dalam suatu urutan. *Multi-head attention* memungkinkan model untuk secara bersamaan memperhatikan informasi dari subruang representasi yang berbeda pada posisi yang berbeda.

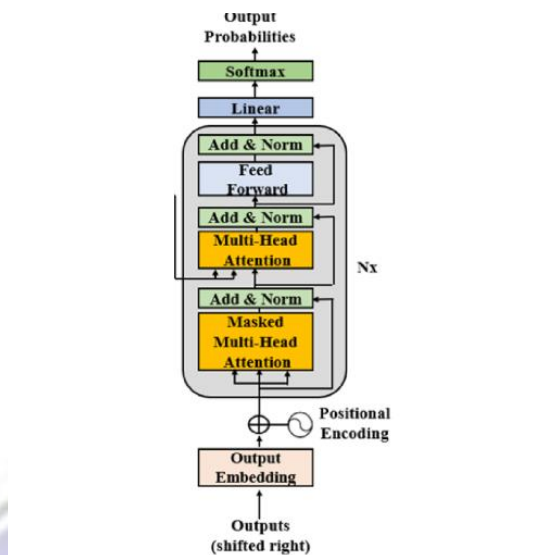
- Sub-layer 2 feed forward *network*

Setelah melewati sublayer multihead attention, *output* nya diteruskan ke feed forward sublayer. Feed forward sublayer

terdiri dari dua lapisan linear yang dipisahkan oleh fungsi aktivasi non-linear seperti ReLU (*Rectified Linear Unit*). Lapisan pertama memproyeksikan *input* dari dimensi model ke dimensi yang lebih besar, kemudian fungsi aktivasi diterapkan, dan lapisan kedua memproyeksikan kembali hasilnya ke dimensi asli. Proses ini membantu dalam meningkatkan kapasitas model untuk menangkap dan mempelajari representasi yang lebih kompleks dari data *input*. Selain itu, *feed forward sublayer* beroperasi secara independen pada setiap posisi dalam urutan *input*, sehingga memungkinkan pemrosesan paralel yang efisien. Kombinasi dari sublayer *multihead attention* dan *feed forward sublayer* memungkinkan *transformer* untuk menjadi model yang sangat efektif dalam berbagai tugas pemrosesan bahasa alami (NLP).

2.2.6.2 Decoderstack

Arsitektur Transformer, yang diperkenalkan oleh Vaswani et al. (2017) dalam makalah mereka "Attention is All You Need," telah menjadi landasan banyak model pemrosesan bahasa alami (NLP) modern[16]. Salah satu komponen utama dari arsitektur *Transformer* adalah *stack decoder*, yang berperan penting dalam menghasilkan *output* selama proses penerjemahan atau penguraian teks.



Gambar 2. 5 Blok *Decoder*[17]

1. Struktur Dasar *Decoder Stack*

Decoder stack pada transformer Gambar 2. 5 Blok *Decoder* terdiri dari beberapa lapisan identik (biasanya enam), masing-masing terdiri dari tiga sublayer utama:

- *Masked Multihead Self-Attention*
- *Multihead Attention* dengan *Encoder-Decoder Attention*
- *Feed Forward Neural Network*

Setiap sublayer ini diikuti oleh mekanisme normalisasi layer (*layer normalization*) dan koneksi residual untuk meningkatkan stabilitas dan efisiensi pelatihan.

2. *Masked Multihead Self-Attention*

Pada sublayer pertama, *masked multihead self-attention*, model memperhatikan semua posisi sebelumnya dalam urutan *output*. *Masking* dilakukan untuk memastikan bahwa prediksi pada posisi tertentu hanya bergantung pada posisi sebelumnya dan tidak pada posisi saat ini atau yang akan datang. Ini memungkinkan model untuk menghasilkan urutan *output* secara *autoregressive*, di mana setiap token dihasilkan satu per satu.

3. *Multihead Attention* dengan *Encoder-Decoder Attention*

Sublayer kedua adalah *multihead attention* yang menghubungkan informasi dari *encoder* ke *decoder*. Di sini, *queries* berasal dari lapisan *decoder* sebelumnya, sedangkan *keys* dan *values* berasal dari *output encoder*. Ini memungkinkan model untuk memperhatikan konteks *input* saat menghasilkan setiap token *output*, memastikan bahwa *output* selaras dengan *input*.

4. *Feed Forward Neural Network*

Sublayer ketiga adalah *feedforward neural network* (FFN), yang terdiri dari dua lapisan linear dengan fungsi aktivasi non-linear di antaranya (biasanya ReLU). FFN diterapkan secara independen pada setiap posisi, dan proses ini membantu dalam mempelajari representasi yang lebih kompleks dari data *input*.

5. Normalisasi Layer dan Koneksi Residual

Setiap sublayer diikuti oleh normalisasi layer (layer normalization) dan koneksi residual. Normalisasi layer membantu dalam mengatasi masalah pelatihan dengan menjaga skala aktivasi tetap stabil. Koneksi residual memungkinkan informasi asli mengalir dengan lebih mudah melalui jaringan, yang membantu dalam mengatasi masalah gradien menghilang (*vanishing gradients*) selama pelatihan.

2.2.7 Bidirectional Encoder Representations from Transformers (BERT)

BERT (*Bidirectional Encoder Representations from Transformers*) adalah algoritma *Natural Language Processing* (NLP) yang dikembangkan oleh Google dan diperkenalkan oleh para peneliti Google AI pada tahun 2018. BERT memberikan hasil yang optimal dalam berbagai tugas NLP seperti *question answering*, *natural language inference*, *classification*, dan *general language understanding evaluation*. BERT hadir dalam dua varian, yaitu BERT-base dan BERT-large. BERT-base memiliki parameter $L=12$, $H=768$, $A=12$, dengan total parameter 110 juta, sedangkan BERT-large

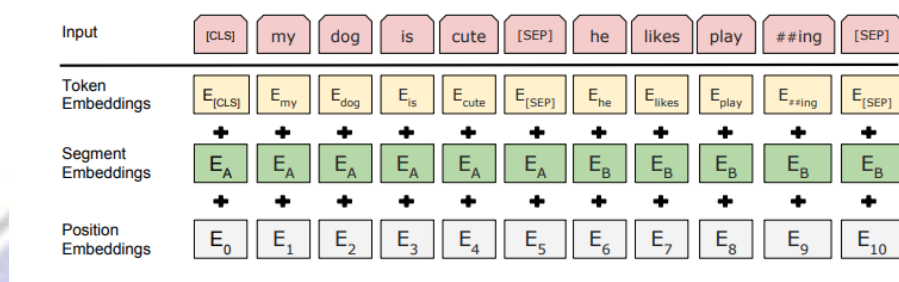
memiliki parameter $L=24$, $H=1024$, $A=16$, dengan total parameter 340 juta.[17]

Input yang diterima oleh BERT adalah vektor angka yang dihasilkan melalui teknik *word embedding* seperti Gambar 2.7. Proses *embedding* ini mengubah setiap token dalam urutan *input* menjadi representasi vektor. Karena arsitektur *Transformer* tidak memiliki koneksi berulang, posisi setiap token dalam urutan *input* harus direpresentasikan secara eksplisit dengan menambahkan vektor *positional encoding* ke *input embedding*. Urutan *input* beserta *positional encoding*-nya kemudian dimasukkan ke mekanisme *multi-head self-attention*. Mekanisme ini memungkinkan model untuk fokus pada bagian-bagian berbeda dalam urutan *input* pada setiap lapisan dan menangkap ketergantungan jarak jauh antar token. Setelah melalui mekanisme *self-attention*, *output* diproses melalui jaringan saraf *feed-forward* yang menerapkan transformasi non-linear pada setiap posisi secara independen. Untuk meningkatkan pelatihan model dan membantu konvergensi lebih cepat, digunakan *residual connections* dan lapisan *normalization*. *Residual connections* memungkinkan aliran gradien lebih mudah melalui jaringan, sedangkan lapisan *normalization* membantu menstabilkan distribusi nilai *output*.

Arsitektur *Transformer* terdiri dari *stack encoder* dan *decoder*. *Stack encoder* bertanggung jawab untuk meng-*encode* urutan *input*, sementara *stack decoder* bertugas menghasilkan urutan *output*. Pada *stack decoder*, mekanisme *self-attention* dimasking sehingga setiap posisi hanya bisa fokus pada posisi sebelumnya dan posisi saat ini. Ini mencegah model melihat token di masa depan selama proses prediksi. Selama proses decoding, *stack decoder* juga memperhatikan *output* dari *stack encoder*, memungkinkan model menggunakan informasi dari urutan *input* untuk menghasilkan urutan *output* yang sesuai [16].

BERT dilatih untuk memahami hubungan kontekstual antara kata-kata dalam sebuah kalimat menggunakan data teks dalam jumlah besar. Khususnya, BERT dilatih pada dua tugas *Masked Language Modeling*

(MLM) dan *Next Sentence Prediction* (NSP). Dalam MLM, beberapa kata dalam sebuah kalimat secara acak diganti dengan token [MASK], dan model harus memprediksi kata aslinya. Tugas ini membantu model memahami makna kata dalam konteks. Dalam NSP, model diberikan dua kalimat dan harus memprediksi apakah kalimat kedua kemungkinan mengikuti kalimat pertama. Tugas ini membantu model memahami hubungan antar kalimat [17]



Gambar 2. 6 Word Embending[20]

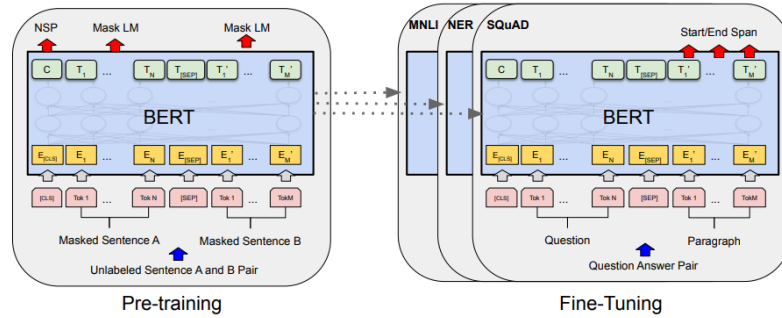
Representasi *input* BERT ditunjukkan pada Gambar 2.6. langkah-langkah tokenisasi dalam BERT adalah sebagai berikut:

1. Tokenisasi: Memecah teks menjadi token-token yang terdiri dari kata-kata. BERT menggunakan tokenisasi *WordPiece*, di mana beberapa token dapat dibagi lagi menjadi sub-token.
2. Token Embeddings: BERT menambahkan dua token khusus ke awal dan akhir setiap kalimat, yaitu [CLS] dan [SEP]. Token [CLS] digunakan untuk merepresentasikan keseluruhan kalimat di awal, sementara token [SEP] digunakan untuk memisahkan kalimat dalam urutan *input*.
3. Konversi Token menjadi ID: Setiap token dalam *input* dikonversi menjadi ID token yang sesuai menggunakan kamus token yang telah ditetapkan. Setiap ID token kemudian dikonversi menjadi vektor dengan mengambil nilai embedding dari matriks embedding kata yang telah dilatih sebelumnya.

4. *Segment Embeddings*: Jika *input* terdiri dari dua kalimat, setiap token dalam *input* diberi tanda sebagai milik kalimat pertama atau kedua dengan memberikan segmen ID 0 atau 1.

5. *Position Embedding*: BERT menggunakan *position embedding* untuk menambahkan informasi posisi absolut ke representasi token dengan menambahkan vektor posisional yang telah ditentukan sebelumnya ke setiap vektor token.

BERT menggunakan dua paradigma pelatihan, yaitu *pre-training* dan *fine-tuning*, yang ditunjukkan pada Gambar 2.7. *pre-training* adalah proses *unsupervised learning* di mana model dilatih pada dataset tanpa label untuk mengekstrak pola. Google melatih model ini pada BooksCorpus (800 juta kata) dan English Wikipedia (2,5 miliar kata). Proses *pre-training* melibatkan dua tugas *masked language modeling* (MLM) dan *next sentence prediction* (NSP). Selama *fine-tuning*, model dilatih kembali pada tugas *downstream* dengan data berlabel, menyesuaikan parameter model BERT yang telah dilatih sebelumnya untuk mengoptimalkan kinerja pada tugas tersebut. *Fine-tuning* melibatkan penambahan lapisan khusus untuk tugas di atas model BERT yang telah dilatih sebelumnya dan melatih seluruh model dari awal hingga akhir pada data tugas tersebut. Lapisan khusus ini memiliki jumlah parameter yang jauh lebih kecil dibandingkan model BERT utama. Selama *fine-tuning*, model dilatih dengan tingkat pembelajaran yang lebih rendah dibandingkan saat *pre-training* karena model sudah belajar fitur umum bahasa dan hanya perlu mempelajari fitur khusus dari tugas *downstream* [18].



Gambar 2. 7 Pre-Training dan Fine Tuning BERT [20]

Setelah melalui beberapa lapisan BERT, *output* terakhir adalah sekumpulan vektor yang mewakili setiap token dalam *input*. Untuk tugas klasifikasi, vektor yang sesuai dengan token khusus [CLS] sering digunakan sebagai representasi dari seluruh *input*. Representasi ini kemudian dikirim ke lapisan klasifikasi yang terdiri dari lapisan linear (*dense layer*) untuk menghasilkan logits.

Rumus perhitungan logits untuk setiap kelas iii diberikan oleh:

$$Logits = X \cdot W + b \quad (2.1)$$

- X adalah vektor representasi akhir dari token [CLS], yang direpresentasikan dengan nilai numerik hasil dari lapisan-lapisan Transformer dalam model seperti BERT.
- W adalah matriks bobot yang dipelajari selama proses pelatihan model.
- b adalah vektor bias yang juga dipelajari selama proses pelatihan.

Logits ini kemudian dilewatkan melalui fungsi aktivasi *softmax* untuk menghasilkan probabilitas prediksi untuk setiap kelas. Fungsi *softmax* mengubah logits menjadi probabilitas dengan rumus:

$$\text{Softmax}(z_j) = \frac{e^{z_j}}{\sum_{j=1}^k e^{z_j}} \quad (2.2)$$

Yang mana:

1. e^{z_j} = nilai eksponensial logits
2. $\sum_{j=1}^k e^{z_j}$ = jumlah keseluruhan eksponen logits
3. e = euler (2.718281828459045)

Fungsi ini memastikan bahwa semua probabilitas hasil prediksi berada dalam rentang $[0,1]$ dan jumlah totalnya adalah 1. Dengan kata lain, fungsi *softmax* menormalisasi nilai logits menjadi distribusi probabilitas yang dapat diinterpretasikan sebagai probabilitas prediksi untuk setiap kelas.

Penjelasan di atas memberikan gambaran tentang bagaimana BERT memproses *input*, menghitung logits, dan menggunakan fungsi *softmax* untuk menghasilkan probabilitas prediksi dalam tugas klasifikasi.

2.2.8 Sentimen Indonesia Lexicon

Sentimen Indonesia lexicon merupakan salah satu pustaka berbahasa Indonesia yang dibangun oleh Koto F dan Rahmaningtyas pada tahun 2017 dengan nama InSet(Indonesia Sentimen) lexicon[19]. Sumber data ini dibangun dengan memuat 3609 kata positif dan 6609 kata negatif dengan pembobotan dari -5 hingga 5. InSet lexicon dibangun dengan memproses data tweet berbahasa Indonesia sekitar 1000 data dengan memberi label setiap kata secara manual berdasarkan polaritasnya kemudian ditingkatkan dengan menambahkan stemming dan set sinonim. [19]. InSet dalam penelitian ini akan dijadikan sumber data dalam proses labelling data.

2.2.9 Confusion Matrix

Confusion matrix merupakan perhitungan untuk mengukur berbagai *performance metrics* terhadap kinerja model yang telah dibangun[2]. Dalam penelitian ini untuk mengukur akurasi dari model yang dibangun menggunakan *confusion matrix* yang merangkum kinerja model pembelajaran mesin pada sekumpulan data uji. Matriks ini digunakan untuk menampilkan jumlah kejadian yang diprediksi dengan benar dan salah oleh model. Matriks ini sangat berguna untuk menilai kinerja model klasifikasi, yang bertujuan untuk memprediksi label kategoris untuk setiap *input*. Matriks konfusi Tabel 2.2 menunjukkan jumlah kejadian yang diprediksi oleh model pada data uji.

- True Positive (TP): Model dengan benar memprediksi titik data sebagai positif.
- True Negative (TN): Model dengan benar memprediksi titik data sebagai negatif.
- False Positive (FP): Model salah memprediksi titik data sebagai positif.
- False Negative (FN): Model salah memprediksi titik data sebagai negatif.

Tabel 2. 2 Confusion Matrix

		Predicted Label	
		Negative	Positive
True Label	Negative	True Negative	False Positive
	Positive	False Negative	True Positive

Berikut metrik evaluasi yang digunakan meliputi

- Akurasi (Accuracy)

Akurasi adalah rasio prediksi yang benar (positif dan negatif) terhadap total jumlah prediksi.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (2.3)$$

TP = True Positive,

TN = True Negative,

FP = False Positive,

FN = False Negative.

- Presisi (Precision)

Presisi adalah rasio true positive terhadap semua prediksi positif yang dilakukan oleh model.

$$\text{Precision} = \frac{TP}{TP+FP} \quad (2.4)$$

- Recall (Sensitivitas atau Recall)

Recall adalah rasio true positive terhadap semua data aktual yang seharusnya positif.

$$\text{Recall} = \frac{TP}{TP+FN} \quad (2.5)$$

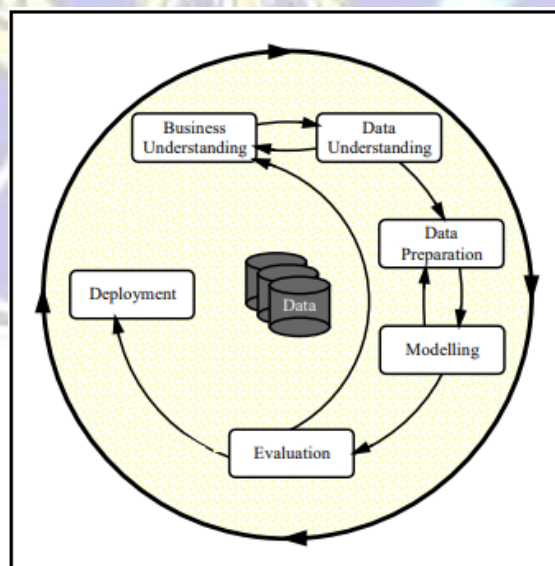
- F1-score

F1-score adalah rata-rata harmonik dari presisi dan recall, memberikan gambaran yang seimbang tentang performa model.

$$F1 - \text{score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (2.6)$$

2.2.10 CRISP-DM (*Cross-Industry Standard Process for Data Mining*)

Data mining membutuhkan standart yang dapat menerjemahkan masalah ke tugas *data mining*, mentransformasikan data dan teknik penambangan data yang tepat, serta menyediakan sarana untuk mengevaluasi efektivitas hasil dan dokumentasi. CRISP-DM atau *Cross-Industry Standard Process for Data Mining* merupakan metode dalam proses data mining dengan tujuan agar proyek penambangan data besar menjadi lebih murah, lebih andal, lebih iterative atau dapat diulang, mudah dikelola, dan lebih cepat[20].



Gambar 2. 8 Proses dalam CRISP-DM[20]

Berdasarkan Gambar 2.8, proses yang dilalui dalam pengerjaan data mining menggunakan metode CRISP-DM dijabarkan seperti berikut.

1. *Bussiness Understanding*

Tahap ini berfokus terhadap pemahaman tujuan dan persyaratan proyek dari perspektif bisnis, dan mengubah pengetahuan menjadi definisi masalah penambangan data, kemudian rencana proyek awal dirancang untuk mencapai tujuan.

2. *Data Understanding*

Tahapan pemahaman data ini dimulai dengan pengumpulan data dilanjutkan dengan eksplorasi data untuk mengidentifikasi masalah kualitas data, menemukan informasi atau untuk mendeteksi subset yang menarik untuk membentuk hipotesis untuk informasi tersembunyi. Ada hubungan erat antara Pemahaman Bisnis dan Pemahaman Data.

3. *Data Preparation*

Proses ini membangun kumpulan data akhir (data yang akan dimasukkan ke dalam alat pemodelan) dari data mentah sebelumnya. *Data preparation* kemungkinan akan dilakukan beberapa kali, dan tidak dalam urutan yang ditentukan. Tugas meliputi beberapa proses seperti pembersihan data, konstruksi atribut baru, serta transformasi data untuk pemodelan.

4. *Modelling*

Pada tahapan ini, berbagai teknik *modelling* dipilih untuk diterapkan, dan parameternya dikalibrasi ke nilai optimal. Biasanya, ada beberapa teknik untuk jenis masalah penambangan data yang sama. Beberapa teknik memerlukan format data tertentu. Tahap ini sangat berhubungan erat dengan data preparing karena mempertimbangkan kualitas data.

5. *Evaluation*

Mengevaluasi model secara lebih menyeluruh, dan meninjau langkah-langkah yang dijalankan untuk membangun model untuk

memastikan model mencapai tujuan bisnis yang direncanakan sudah sesuai.

6. Deployment

Pengetahuan yang diperoleh perlu diatur dan disajikan agar dapat menjadi insight bagi orang lain. Tahap ini dapat dengan membuat laporan atau membuat visualisasi hasil data mining.

2.2.11 Python

Python merupakan bahasa pemrograman populer yang sering digunakan dalam berbagai pengembangan seperti *data mining*, *machine learning* hingga *web development*. Python memiliki struktur data tingkat tinggi yang efisien dan pendekatan sederhana namun efektif terhadap pemrograman berorientasi objek. Sintaksnya yang elegan dan pengetikan dinamis, bersama dengan sifatnya yang terinterpretasi, menjadikannya bahasa yang ideal untuk *scripting* dan pengembangan aplikasi cepat di berbagai bidang pada sebagian besar *platform*[21].

2.2.12 Flask

Flask adalah kerangka kerja mikro berbasis Python yang dirancang agar mudah dipahami dan digunakan, menawarkan fleksibilitas melalui ekstensi tanpa menambah kompleksitas yang tidak perlu. Dengan tiga dependensi utama, yaitu Werkzeug, Jinja2, dan Click, Flask menyediakan inti yang kuat untuk pengembangan web[22]. Meski tidak mendukung fitur tingkat lanjut secara bawaan, seperti akses basis data atau autentikasi pengguna, fitur-fitur ini dapat ditambahkan melalui ekstensi, memberi pengembang kebebasan untuk memilih alat yang sesuai dengan kebutuhan proyek. Flask sangat cocok bagi mereka yang menginginkan kerangka kerja Python yang ringan namun tetap fungsional dan dapat diperluas.

2.2.13 HTML

HTML atau *Hypertext Markup Language* merupakan bahasa *markup* yang menjadi standarisasi pengembangan halaman website yang dapat ditampilkan pada *web browser*. HTML muncul pertama kali pada 1989 oleh Tim Berners Lee yang kemudian dikembangkan oleh W3C.[23]

2.2.14 CSS

Cascading style sheets atau dikenal dengan CSS merupakan bahasa untuk mengatur tampilan atau *layout website* bersanding dengan HTML dalam penggunaannya. Bahasa ini mendukung berbagai bahasa *markup* seperti HTML, XHTML, XML, SVG (*Scalable Vector Graphics*) dan Mozilla XUL (*XML User Interface Language*).[23]

2.2.15 Black Box Testing

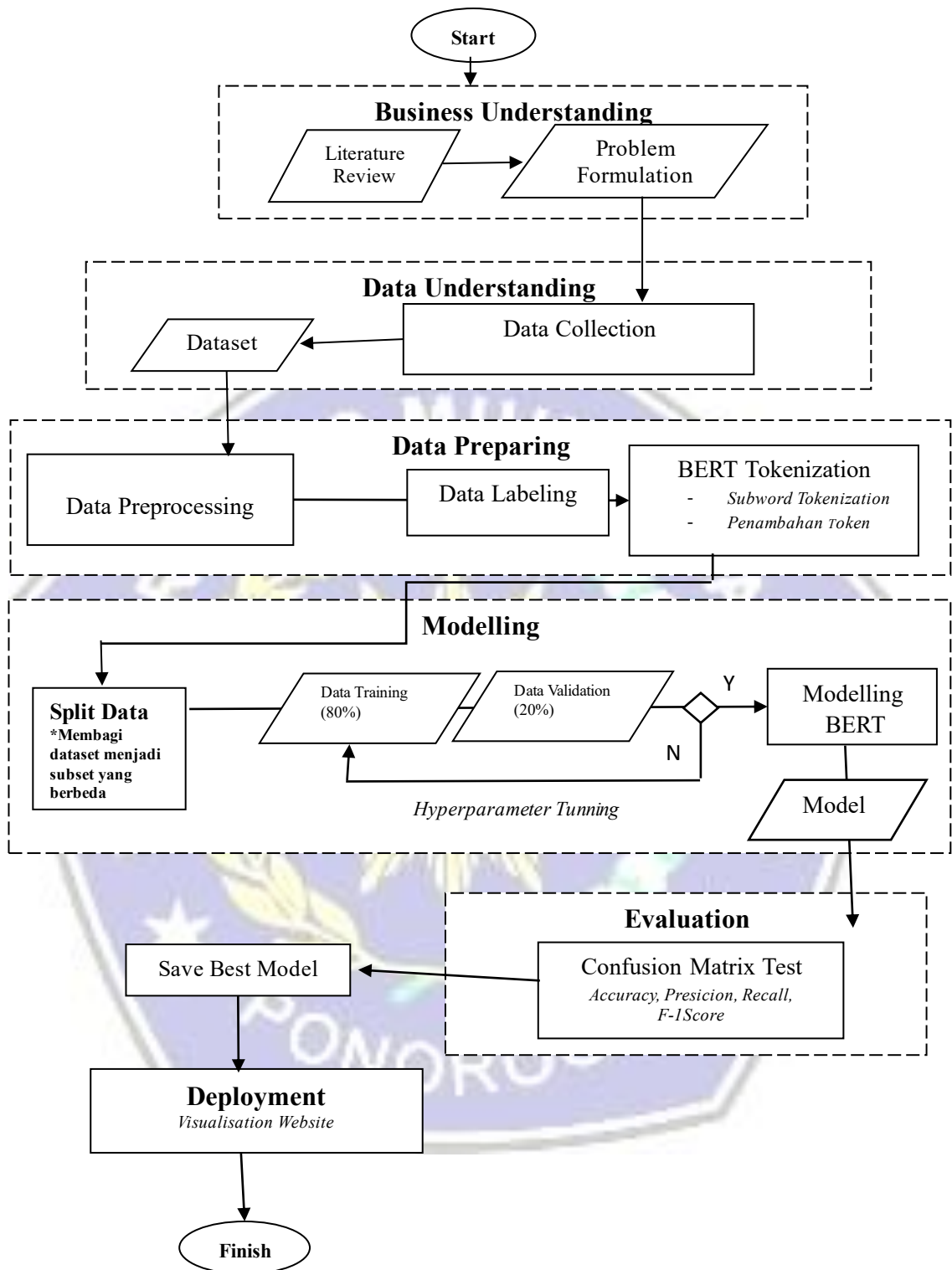
Dalam pengembangan aplikasi ataupun website diperlukan tindakan untuk melihat hasil yang sudah dirancang dengan pengujian atau testing. *Black Box* merupakan salah satu teknik pengujian dengan fokus terkait spesifikasi fungsional dari perangkat lunak yang telah dibangun. [24]

BAB III

METODE PENELITIAN

Dalam penelitian ini menerapkan metodologi CRISP-DM (*Cross-Industry Standard Process for Data Mining*). Menerapkan CRISP-DM dalam proyek analisis sentimen untuk penerapan *E-Parking* "PARKIR-GO" oleh Pemerintah Kabupaten Ponorogo dengan algoritma BERT menjadi pilihan peneliti karena menyediakan kerangka kerja terstruktur dan iteratif yang memastikan setiap tahap proyek dilaksanakan dengan hati-hati dan sistematis. CRISP-DM dimulai dengan pemahaman bisnis yang mendalam untuk memastikan relevansi hasil analisis, dilanjutkan dengan persiapan data yang komprehensif untuk menangani data tidak terstruktur dari media sosial.

Tahap pemodelan memungkinkan penggunaan dan fine-tuning BERT untuk analisis sentimen yang akurat, sementara evaluasi menyeluruh memastikan model memenuhi kriteria performa yang diinginkan. Pada tahap terakhir adalah tahap *deployment* dan *monitoring* memastikan model dapat digunakan secara efektif dan diperbarui sesuai kebutuhan, memberikan wawasan berharga bagi keputusan pemerintah. Kerangka kerja yang akan dilaksanakan dijelaskan pada Gambar 3.1 Diagram Alur Metodologi CRISP-DM sebagai berikut.



Gambar 3. 1 Diagram Alur Metodologi Penelitian

3.1. Business Understanding

Pada tahapan ini dilakukan beberapa langkah untuk mendukung dan menjadi salah satu fondasi penelitian. Beberapa tahapan yang dilalui sebagai berikut:

- *Literature Review* (Tinjauan Pustaka)

Peneliti melakukan pengumpulan data dan informasi terkait berdasarkan jurnal ilmiah, buku, artikel, website dan sumber akurat lainnya. Setelah mengumpulkan beberapa referensi terkait dilakukan pengembangan ide berdasarkan masalah yang diamati di sosial media terkait penerapan *E-Parking* Kabupaten Ponorogo dan melakukan perbandingan untuk mendapatkan landasan teori pada penelitian. Seperti yang telah di paparkan pada Tabel 2.1 literatur yang dijadikan landasan memiliki topik serupa yakni sentimen analisis dengan metode NLP terkhusus metode BERT di berbagai macam sektor.

- *Problem Formulation* (Identifikasi Masalah)

Setelah mengumpulkan landasan teori dan informasi dari beberapa sumber, peneliti kemudian merumuskan permasalahan yang dijadikan landasan penelitian ini dengan mempelajari penelitian terdahulu Tabel 2.1. yang berkaitan dengan BERT. Penelitian Rhini Fatmasaria,dkk(2023) dilakukan analisa terkait komentar sosial media twitter dan tiktok sebanyak 658 data hampir sama dengan jumlah sample yang dimiliki pada penelitian ini yakni sekitar 450 data menggunakan beberapa teknik yang diuji menunjukkan hasil BERT dengan akurasi sebesar 90% lebih tinggi dibanding yang lain[3]. Penelitian Nanang H (2023) menunjukkan algoritma BERT mencapai akurasi *training* sebesar 93% dan akurasi testing sebesar 92%, dengan marco avg dari f1 score sebesar 92%[2]. Dengan demikian, algoritma BERT terbukti efektif dalam mengklasifikasikan teks artikel berita, terutama dalam kasus dataset yang cukup besar dan tidak seimbang dibanding metode random forest yang memiliki akurasi *training* 81% dan Naïve Baiyes 78%. Kesimpulan dari beberapa penelitian tersebut menunjukan kemampuan yang mumpuni dari BERT dalam menganalisis

teks yang kompleks dalam analisis sentimen. Penelitian ini merumuskan masalah pada model BERT dengan hasil akhir sentimen analisis penerapan *E-Parking* Kabupaten Ponorogo dan tingkat akurasi dari model ini. Data pada penelitian ini adalah komentar sosial media dengan topik *E-Parking* Ponorogo pada jangka waktu januari-mei 2023. Pada penelitian ini dibutuhkan perencanaan kebutuhan seperti berikut.

a. Waktu Penelitian

Penelitian akan melewati beberapa proses sehingga perlu ditetapkan waktupenelitian diterangkan pada Tabel 3.1

Tabel 3. 1 Waktu Penelitian

Tahapan	Tools	Objek
Data Collection	Apify IGCommentExport	Konten dengan topik E- Parking
Data Preparation	Python	Dataset
Modeling	Python	Dataset

b. Perangkat Penelitian

Perangkat yang digunakan dalam proses pengujian penelitian ini sebagai berikut:

- 1) Data dikumpulkan dan diolah dengan menggunakan laptop berspesifikasi berikut:
 - Processor : Intel Core i3 10th Gen
 - RAM : 4 GB,
 - Storage : SSD 163 GB
- 2) Perangkat diatas dapat memanfaatkan layanan cloud gratis Google yakni Google Colaboratory dengan GPU yang tersedia sebagai akselerator perangkat.
- 3) Perangkat lunak yang digunakan dalam penelitian ini yakni sistem operasi window 11 64-bit, google colab, python, micrisoft edge.

3.2. Data Understanding

Setelah melalui tahapan bussines understanding langkah yang selanjutnya adalah data understanding atau memahami data yang akan digunakan. Pengumpulan data yang dilakukan menerapkan teknik web-scraping otomatis dengan tools IGCommentsExport dan Apify dari postingan akun terkait Tabel 3.2. Selanjutnya hasil scraping disusun menjadi dataset dengan format .csv ditampilkan pada Gambar 3.2 yang kemudian akan memasuki proses *data preparation* dengan *library* python seperti pandas. Berikut rincian pengumpulan data dari sosial media.

User Id	Username	Comment Id	Comment Text	Profile URL	Avatar URL	Date
3211367537	tatik_mujiarti	1.80731E+16	kalo bisa di persulit kenapa harus di pe	https://www.instagram.com/tatik	https://scontent-cgk1-2.5/23/2023, 12:01:10 PM	
47616636524	ramangabdulazis	1.7994E+16	Kampungan ☺	https://www.instagram.com/raman	https://instagram.flir5-1.5/23/2023, 12:03:20 PM	
44103852692	marissasr	1.78623E+16	@alistaavs bos berek	https://www.instagram.com/maris	https://scontent-cgk1-2.5/23/2023, 12:35:04 PM	
1808322242	andika.dbh	1.80654E+16	Teknologi melebihi akal sehat, 2023 n	https://www.instagram.com/andika	https://scontent-cgk1-2.5/23/2023, 12:42:25 PM	
2135889890	ragiel_mangku	1.79938E+16	Mugo ae nek khilangan di tmpt parkir	https://www.instagram.com/ragiel	https://scontent-cgk1-2.5/23/2023, 12:55:26 PM	
9180409937	arin9499	1.79993E+16	Maleh gk minat dolan rmo min . @pon	https://www.instagram.com/arin94	https://scontent-cgk1-2.5/23/2023, 12:57:14 PM	
43635174987	herysantu	1.80119E+16	Ngrusi urusan sing gak urgent	https://www.instagram.com/herys	https://scontent-cgk1-2.5/23/2023, 1:01:00 PM	
5923339307	nadflx.jv	1.79666E+16	r ndwe e money	https://www.instagram.com/nadflx	https://scontent-cgk1-2.5/23/2023, 1:21:19 PM	
7180642594	optik_tasikmalay	1.79891E+16	Lama lama kembali lagi ke manual	https://www.instagram.com/optik	https://scontent-cgk1-2.5/23/2023, 1:45:06 PM	
1654627163	erikwahyudy	1.78854E+16	Semangatt ☺	https://www.instagram.com/erikw	https://scontent-cgk1-2.5/23/2023, 1:45:26 PM	
1781213145	eline_xiandhia	1.79868E+16	@yogiiprampito nah tu dia.... 📢📢	https://www.instagram.com/eline	https://scontent-cgk1-2.5/23/2023, 2:39:51 PM	
1445963182	hanifahlayli	1.80127E+16	Aku luweh ikhlas duit sewu rongewu n	https://www.instagram.com/hanifa	https://scontent-cgk1-2.5/23/2023, 2:43:30 PM	
2066042336	zazux_yeye	1.7967E+16	Pendapatku tetep penak parkir gratis.	https://www.instagram.com/zazux	https://scontent-cgk1-2.5/23/2023, 3:08:36 PM	
52763097857	dikaambarani	1.7956E+16	Alhamdulillah Semangat Gak Belum B	https://www.instagram.com/dika	https://scontent-cgk1-2.5/23/2023, 4:07:39 PM	
57408281609	adityaanakbaik	1.80601E+16	Nyuen nyuen i lakon	https://www.instagram.com/aditya	https://scontent-cgk1-2.5/23/2023, 4:17:01 PM	
6185169121	humairasyaharan	1.79826E+16	Pendapat ku apik lor mergo ben ora ha	https://www.instagram.com/huma	https://scontent-cgk1-2.5/23/2023, 4:18:37 PM	

Gambar 3. 2 Dataset Hasil Scrapping

Tabel 3. 2 Data Collection

Sosial Media	Akun	Jumlah Konten terkait	Komentar Diambil
Instagram	@infoponorgo	2	200
	@ponorogoupdate	3	267
	@ponorogopictures	1	17
Facebook	Gemasurya	2	22
Youtube	KompasTV	1	32
Total Data Komentar			538

3.3. Data Preparation

Tahapan ini merupakan tahap penting dalam penelitian karena mengubah data yang mentah menjadi data yang siap untuk dipakai dalam tahap *modelling*.

3.3.1 Data Preprocessing

Proses pembersihan data dilakukan agar dataset yang dimiliki menjadi data yang bersih dan akurat untuk proses analisa. Data yang bersih dan berkualitas akan mendukung pengambilan keputusan yang lebih baik dan akurat. Beberapa proses dilalui untuk mendapatkan data yang bersih sebagai bahan modeling di tahap selanjutnya.

1. Pembersihan Data Duplikat

Proses ini menghilangkan data yang sama persis karena proses penggabungan data saat berada di tahap scraping. Metode yang dipakai yakni `drop_duplicates()` dengan python. Setelah itu dilakukan pemilihan kolom yang akan digunakan yakni kolom text karena berisi data komentar.

2. Pembersihan simbol dan *noise* (#, @, emoticon, simbol non alfabetik, dan spasi awal dan akhir kalimat)

Proses ini memanfaatkan modul *regular expression*(re) yang dimiliki python. Modul ini menyediakan fungsi-fungsi yang mendukung penggunaan regular expression dalam pengolahan teks. Berikut contoh hasil proses ini.

Tabel 3. 3 Data Cleansing

Sebelum	Sesudah
text 0 Sip...buat ngurangi pungli 1 @ricardo_kneff kalau di jakarta e parking bias... 2 Bayar pakek gopay,dana opo sopipay? 3 Pembayaran paling gampang Jane gae Qris 4 Kalau gak di mulai tdk akan pernah mulai.. bis... ... 441 Pokok aku padamu pak/bu semerintah, arep digaw... 442 Siji d kon ngno... Liane do ngaleh...	text 0 sipbuat ngurangi pungli 1 kalau di jakarta e parking biasanya pakai e mo... 2 bayar pakek gopaydana opo sopipay 3 pembayaran paling gampang jane gae qris 4 kalau gak di mulai tdk akan pernah mulai bismi... ... 441 pokok aku padamu pakbu semerintah arep digawe ... 442 siji d kon ngno liane do ngaleh

3. Case Folding

Setelah melewati data cleansing seperti Tabel 3.3, proses pengubahan huruf menjadi huruf kecil untuk membantu dalam membuat

teks lebih konsisten dengan menghilangkan variasi yang disebabkan oleh perbedaan kapitalisasi. Misalnya, kata "Hello" dan "hello" dianggap sama setelah *lowercase*, yang mengurangi kompleksitas dan ukuran kosakata yang perlu dipelajari oleh model. Proses ini juga menggunakan metode *text.lower()* untuk mengubah semua karakter dalam *string text* menjadi huruf kecil atau *lowercase* seperti pada Tabel 3.4 dibawah ini.

Tabel 3. 4 Case Folding

Sebelum diproses	Lowercase
Program Yang Harus Didukung Karena Positif	program yang harus didukung karena positif

4. Slang word Removal

Data yang sebelumnya kemudian di proses untuk menghilangkan kata-kata slang yang ada. Kamus yang dipakai dalam proses ini berasal dari <https://github.com/nasalsabila/kamus-alay.git> yang mengandung 3592 kumpulan kata.

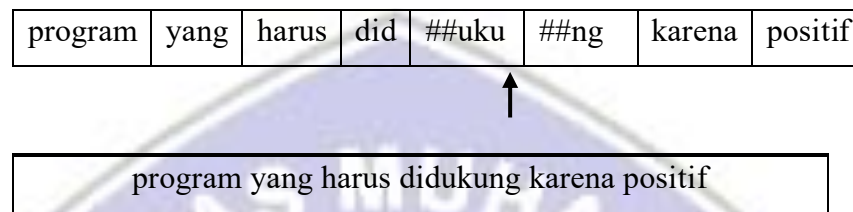
3.3.2 Data Labelling

Tahap *labelling* data menerapkan metode *lexicon* dengan kamus sentimen bahasa indonesia InSet dengan memberi bobot -5 hingga 5. Selanjutnya dikategorikan menjadi tiga sentimen yakni apabila bobot dibawah 0 maka negatif, lebih dari 0 hingga 5 positif, dan lainnya dilabeli dengan sentimen netral.

3.3.3 BERT Tokenization

Langkah dalam representasi *input* di BERT sebagai berikut:

1. Tokenisasi *Wordpiece* menjadi subkata. Beberapa sub kata dapat diawali dengan simbol *##* sebagai penanda token tersebut adalah sufiks dan diikuti subkata lain.

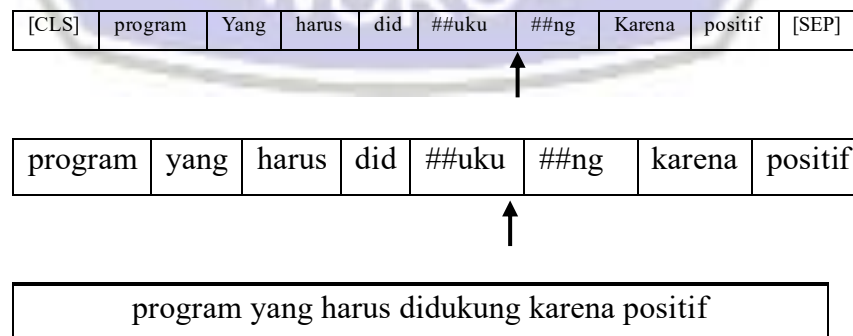


Gambar 3. 3 Token WordPiece

Seperti pada **Gambar 3.3** kalimat dibagi menjadi sub kata yang awalnya “program yang harus didukung karena positif” menjadi “program” “yang” “harus” “did” “##uku” “##ng” “karena” “positif”. Kata didukung dipisah menjadi “did” “##uku” “##ng” karena tidak terdapat dalam *vocabulary* dari model *pre-trained bert-base-multilingual*.

2. Token Embedding

Setiap kalimat akan diberikan token khusus berupa [CLS] diawal dan akan digunakan pada proses klasifikasi sentimen karena token ini melakukan pengumpulan rata-rata token untuk mendapat vektor dari kalimat. Kemudian token [SEP] sebagai pemisah dengan kalimat selanjutnya.



Gambar 3. 4 Token Embedding

Pada Gambar 3.4. merupakan proses *embedding* dengan penambahan token pada kalimat yang sudah dipisah menjadi sub kata sesuai BERT *tokenize*.

3. Token Padding

Karena dalam BERT panjang *input* harus dibuat sama, token [PAD] dapat difungsikan sebagai penambah *input* agar sesuai panjang yang ditentukan. Misalkan panjang maksimal 125 maka pada kalimat contoh membutuhkan token [PAD] hingga memenuhi panjangnya.



Gambar 3. 5 Token Padding

Proses ini dilakukan untuk penyesuaian panjang *input* , contoh pada Gambar 3.5 Menambahkan token *padding* [PAD] untuk menambah panjang token.

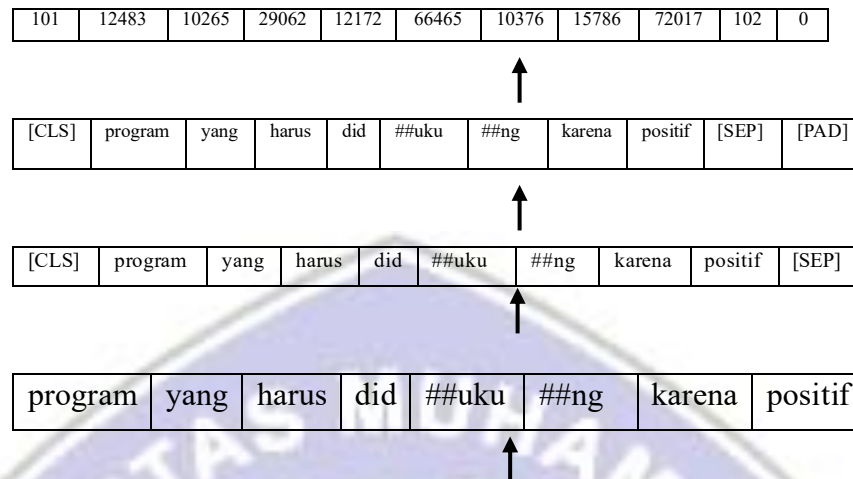
4. Subtitusi ID

Proses ini didapatkan dari indeks kata dalam *vocabulary* pada model bert-base-*multilingual* seperti Gambar 3.6, misalkan token [UNK] memiliki id 100, [CLS] memiliki id 101, [SEP] memiliki Id 102 dan [PAD] memiliki id 0.

[UNK]:	100
[SEP]:	102
[CLS]:	101
[PAD]:	0

Gambar 3. 6 Indeks ID token BERT

Dan untuk kata lain yang muncul akan mendapatkan id yang telah dimuat dalam vocabulary model.



Gambar 3. 7 Tokenisasi ID BERT

Proses pada Gambar 3.7. mengubah token menjadi token numerik yang kemudian akan dipakai pada proses *input* klasifikasi *sentimen* BERT.

5. Sentence Embedding

Sebagai pembeda kalimat satu dengan lainnya proses ini akan memberi angka penanda yang sama di setiap sub katanya seperti Gambar 3.8.

[CLS]	program	yang	harus	did	##uku	##ng	karena	positif	[SEP]	[PAD]
1	1	1	1	1	1	1	1	1	1	0

Gambar 3. 8 Sentence Embedding

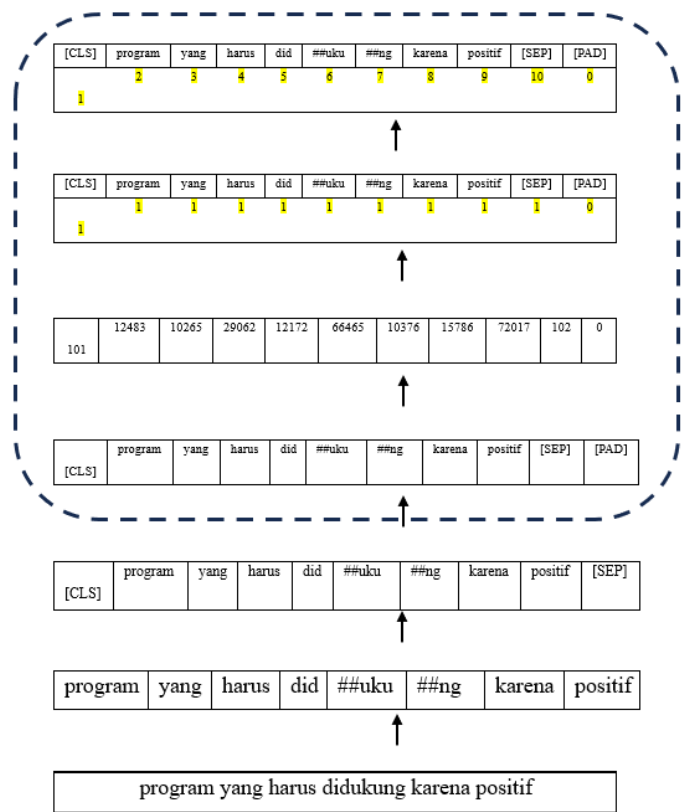
6. Positional embedding

Di dalam BERT diperlukan penomoran posisi tiap kata dalam kalimat agar dapat memahami makna dari kalimat. Pada Gambar 3. 9. menjelaskan tahap penomoran posisi masing-masing token dalam kalimat.

[CLS]	program	yang	harus	did	##uku	##ng	karena	positif	[SEP]	[PAD]
1	2	3	4	5	6	7	8	9	10	0

Gambar 3. 9 Positional Embedding

Pada Gambar 3.10. menunjukkan proses tokenisasi sebagai proses representasi *input*.



Gambar 3. 10 Proses Representasi *Input*

3.4. Modeling

3.4.1 Data Split

Data yang sudah dibersihkan dibagi menjadi data *train* untuk pelatihan model kemudian data *validation* untuk evaluasi. Dikarenakan data yang dimiliki cenderung kecil maka proporsi data train lebih besar dengan perbandingan 80:20. Sebesar 80 persen data asli digunakan sebagai data train dan 20 persen sebagai data validasi. Pembagian kelas pada data *train* yakni 163 sampel kelas negatif(0), 156 sampel kelas positif(2) dan 100 sampel kelas netral(1). Pembagian kelas data validasi yakni 41 sampel kelas negatif(0), 38 sampel kelas positif(2), 26 sampel kelas netral(1).

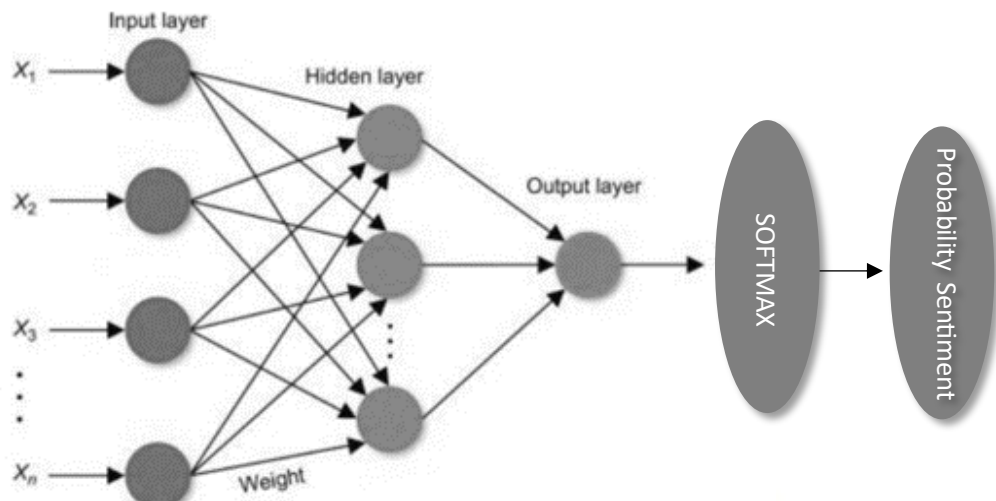
3.4.2 Load Pre-Training & Fine-Tuning BERT model

Pada BERT pelatihan tidak dilakukan dari awal tetapi memanfaatkan model yang telah dilatih sebelumnya(pre-trained) dengan data besar kemudian disesuaikan dengan sedikit pembelajaran untuk tugas baru atau disebut dengan teknik *fine-tune*.

3.4.3 Klasifikasi BERT

Pada model BERT ada beberapa proses yang dilalui mulai dari *input* representation untuk menyiapkan *input* an ke model hingga klasifikasi. Peneliti pada penelitian ini akan menerapkan pre-trained bert-base-multilingual yang kemudian di *fine tuning*. Sebelumnya di tahap *data preparation* data sudah melalui proses sebagai representasi *input* model. BERT memiliki panjang maksimum kalimat *input* yakni 512 karena *encoder transformer* hanya mampu menghasilkan *output* dengan dimensi 512.

Setelah menjadi bentuk representasi *input* pada Gambar 3. 10 selanjutnya akan berlanjut ke proses klasifikasi sentimen. Pada layer model *output* akan menghasilkan vektor dari [CLS] yakni berupa logits. Logits ini yang akan diubah menjadi probability dengan metode aktivasi softmax. Probabilitas dengan *softmax* ini ketika dijumlah keseluruhan akan menjadi tepat 1 dan nilainya diantara 0 hingga 1.



Gambar 3. 11 Tahap Klasifikasi

Dari ilustrasi pada Gambar 3.11 *output* selanjutnya di proses dengan *softmax*. Berikut tahapan memperoleh nilai probabilitasnya:

- Menghitung eksponen elemen logits.

Logits diperoleh dengan melakukan operasi linear pada *output* dari model sebelum diterapkan fungsi aktivasi. Secara matematis, ini dapat diwakili dengan persamaan (2.1)

Misalkan hasil tokenisasi dan embedding untuk kalimat

"program yang harus didukung karena positif"

menghasilkan vektor representasi akhir X . Matriks bobot W akan memiliki dimensi yang sesuai dengan representasi akhir token [CLS] (misalnya, jika menggunakan BERT base, dimensi 768×3 untuk tiga kelas). Vektor bias b adalah vektor dengan panjang yang sama dengan jumlah kelas (misalnya, 3 untuk tiga kelas: positif, netral, negatif). Misalkan nilainya seperti berikut.

$$X = [0.5, -0.3, 0.7, \dots]$$

$$W = [[w_{11}, w_{12}, w_{13}], [w_{21}, w_{22}, w_{23}], \dots, [w_{768}, w_{768}, w_{768}]]$$

$$b = [0.1, -0.2, 0.3]$$

Kemudian dihitung dengan persamaan (2.1).

$$\text{logit1} = X[0].W[0][0] + X[1].W[1][0] + \dots + X[767].W[767][0] + b[0]$$

$$\text{logit2} = X[0].W[0][1] + X[1].W[1][1] + \dots + X[767].W[767][1] + b[1]$$

$$\text{logit3} = X[0].W[0][2] + X[1].W[1][2] + \dots + X[767].W[767][2] + b[2]$$

Hasil logits adalah [logit1, logit2, logit3] atau $[Z_1, Z_2, Z_3]$

Misal hasil diatas adalah $Z = \begin{pmatrix} Z_1 \\ Z_2 \\ Z_3 \end{pmatrix} = \begin{pmatrix} 5 \\ 3 \\ 2 \end{pmatrix}.$

Selanjutnya untuk mendapatkan probabilitas langkah awalnya dengan menghitung eksponensial setiap elemen logits.

$$e^{Z_1} = e^5 = 148.41$$

$$e^{Z_2} = e^3 = 20.09$$

$$e^{Z_3} = e^2 = 7.39$$

- Menjumlahkan seluruh eksponen

$$\sum_{j=1}^k e^{Z_j} = e^5 + e^3 + e^2 = 148.41 + 20.09 + 7.39 = 175.89$$

- Mendapat probabilitas dengan *softmax* dengan persamaan (2.2)

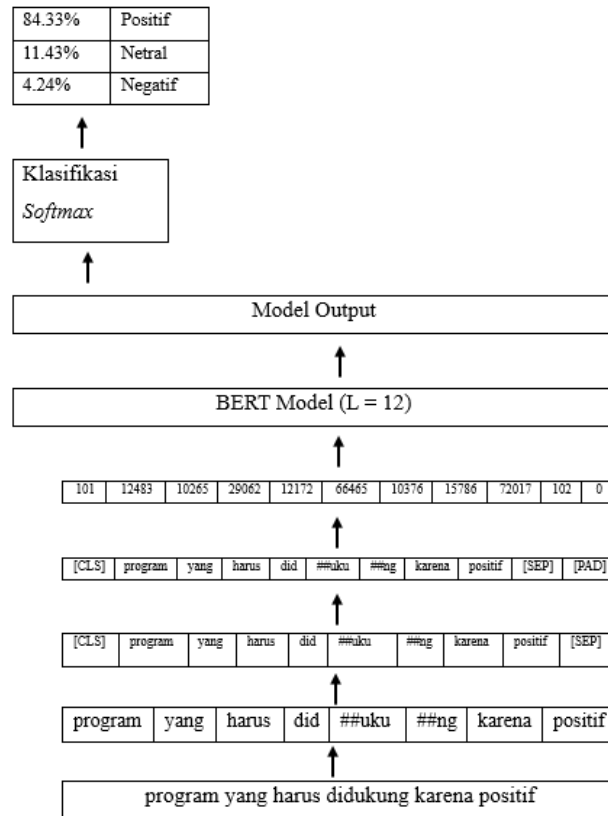
$$\text{Softmax}(z_1) = \frac{e^5}{\sum_{j=1}^k e^{Z_j}} = \frac{148.41}{175.89} = 0.8433$$

$$\text{Softmax}(z_2) = \frac{e^3}{\sum_{j=1}^k e^{Z_j}} = \frac{20.09}{175.89} = 0.1143$$

$$\text{Softmax}(z_3) = \frac{e^2}{\sum_{j=1}^k e^{Z_j}} = \frac{7.39}{175.89} = 0.0424$$

$$\text{Maka hasil probabilitas total } 0.8433 + 0.1143 + 0.0424 = 1$$

Pada contoh penerapan diatas maka *sentimen* dari kalimat tersebut dominan di kelas positif sebesar 0.8433. Gambar 3.12 menunjukkan ilustrasi klasifikasi sentimen BERT secara keseluruhan.



Gambar 3. 12 Ilustrasi Klasifikasi Sentimen BERT

Mekanisme di tahap modelling BERT Gambar 3.12 merupakan lapisan *encoder* berjumlah 12 karena pada penelitian ini menggunakan BERT base yang memiliki jumlah layer tersebut. Pada setiap layer *encoder* melakukan mekanisme *self-attention* seperti contoh mekanisme Gambar 2.3. Setelahnya di lapisan terakhir menghasilkan *output* untuk representasi teks dalam token [CLS] yang kemudian akan dilakukan klasifikasi dengan fungsi aktivasi *softmax* persamaan (2.2) untuk melihat probabilitas sentimen akhir masuk kelas apa.

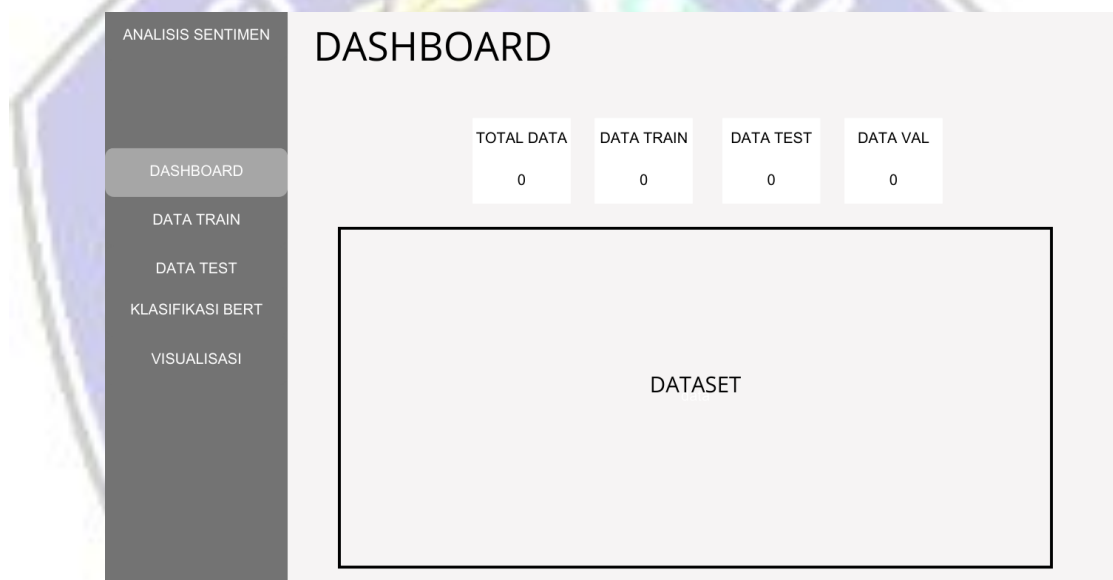
3.5. Evaluation

Evaluasi akurasi, presisi, *recall* dan F1-score adalah langkah penting dalam mengevaluasi performa model klasifikasi, terutama dalam konteks pemrosesan bahasa alami (NLP). Dengan menggunakan persamaan (2.3)

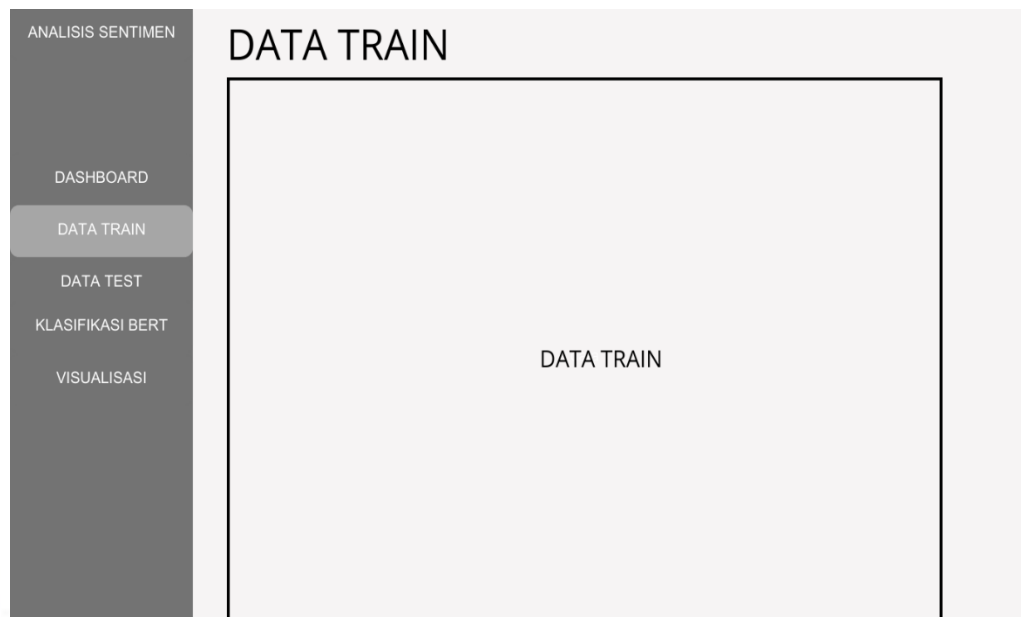
untuk akurasi, persamaan (2.4) untuk presisi, *recall* dengan persamaan (2.5) dan menghitung skor F1 dengan persamaan (2.6)

3.6. Deployment

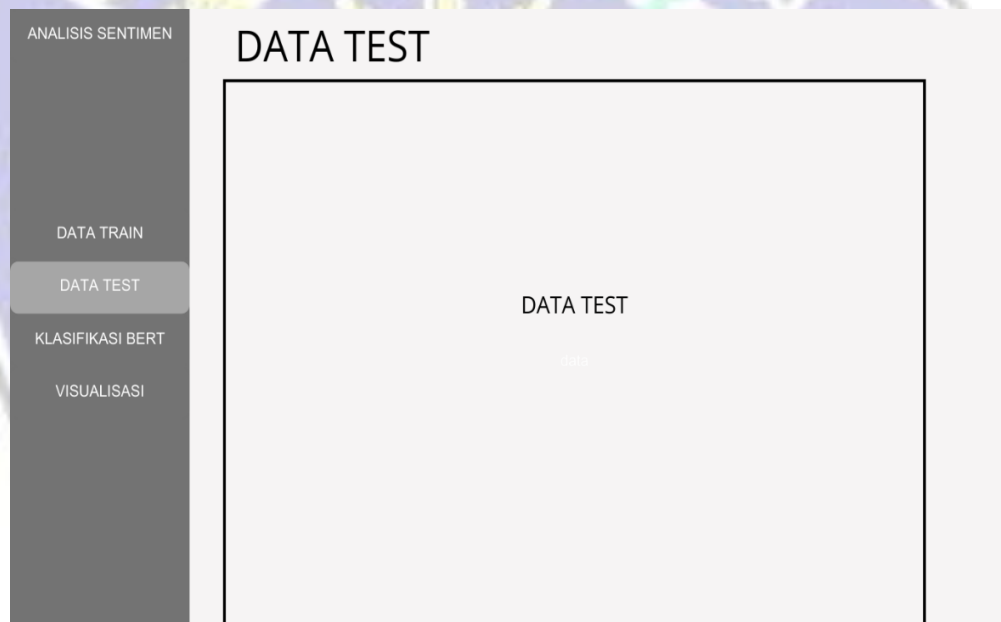
Tahap akhir dengan mengambil Kesimpulan dari hasil pemodelan menggunakan BERT. Selain merancang GUI untuk melihat kemampuan analisis *sentimen* model disusun juga dashboard hasil analisis sentimen *E-Parking* berbasis website sederhana menggunakan *framework* Flask Python sebagai server side, library *pandas* untuk membaca dataset, kemudian HTML, CSS sebagai pondasi tampilan web. Berikut dashboard untuk menampilkan hasil analisis sentimen dengan pendekatan BERT pada Gambar 3.13., Gambar 3.14., Gambar 3.15., Gambar 3.16. dan Gambar 3.17



Gambar 3. 13 Desain Halaman Dashboard



Gambar 3. 14 Desain Halaman Data Train



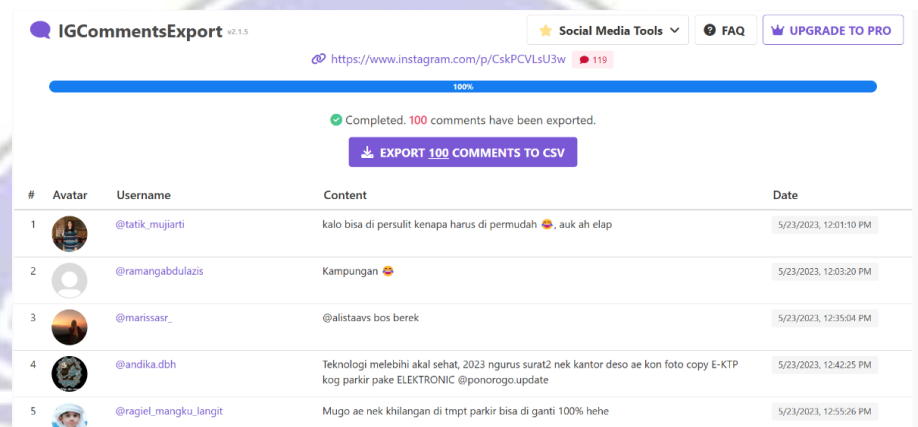
Gambar 3. 15 Desain Halaman Data Test

BAB IV

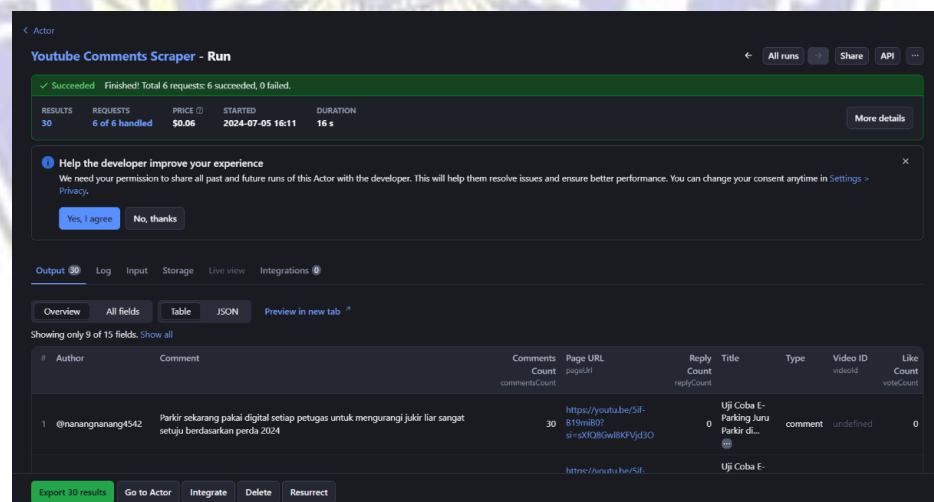
HASIL DAN PEMBAHASAN

4.1. Data Collecting

Data yang digunakan pada penelitian ini merupakan data komentar sosial media yang dikumpulkan dengan metode web scraping menggunakan tools otomatis yaitu apify dan IGComment scraper yang dapat diakses langsung melalui web browser. Berikut adalah proses data collecting yang dilakukan pada Gambar 4.1. dan Gambar 4.2.



Gambar 4. 1 Scrapping Komentar Instagram



Gambar 4. 2 Scrapping Komentar Youtube

1. Data Cleansing

Proses pembersihan data mentah menjadi data bersih agar meningkatkan akurasi dan kinerja model yang akan dijalankan. Langkah pertama yang dilakukan untuk mendapatkan data yang bersih pada penelitian ini adalah menghilangkan null data atau data kosong pada dataset yang dimiliki. Pada pembersihan null data menggunakan metode dropna seperti Gambar 4.5.

```
[ ] #cek null data
df.info()

<class 'pandas.core.frame.DataFrame'>
RangeIndex: 527 entries, 0 to 526
Data columns (total 1 columns):
 #   Column  Non-Null Count  Dtype
---  --
 0   Comment  527 non-null      object
dtypes: object(1)
memory usage: 4.2+ KB

[ ] df.isnull().sum()

Comment    0
dtype: int64

[ ] #MULAI DATA HANDLING
df.dropna(inplace=True)

[ ] df.info()

<class 'pandas.core.frame.DataFrame'>
RangeIndex: 527 entries, 0 to 526
Data columns (total 1 columns):
 #   Column  Non-Null Count  Dtype
---  --
 0   Comment  527 non-null      object
dtypes: object(1)
memory usage: 4.2+ KB
```

Gambar 4. 5 Missing Value Handling

Selanjutnya adalah membersihkan simbol-simbol pada komentar termasuk mention, hastag dan emoticon dilakukan pada proses Gambar 4.6.

```
[ ] import re

# Fungsi untuk menghilangkan mention dan emotikon
def remove_mentions_and_emoticons(text):
    # Menghilangkan mention
    text = re.sub(r'@[^\s]+', '', text)

    # Menghilangkan emotikon
    text = re.sub(r'[\U00010000-\U0010ffff]', '', text) # Menghapus karakter astral, termasuk emotikon

    return text

df['Comment'] = df['Comment'].apply(remove_mentions_and_emoticons)
df['Comment'] = df['Comment'].apply(remove_mentions_and_emoticons)
```

Gambar 4. 6 Pembersihan Simbol

Berikut hasil yang diperoleh dari proses ini pada Gambar 4.7.

```
#Kolom sebelum di proses
df1[['comment']].head()
```

	comment
0	@ricardo_kneff kalau di jakarta e parking bias...
1	Bayar pakek gopay,dana opo sopipay?
2	Pembayaran paling gampang Jane gae Qris??
3	Kalau gak di mulai tdk akan pernah mulai.. bis...
4	Wkwk brharap parkir motor di seluruh Ponorogo...

```
[ ] df.head()
```

	Comment
0	kalau di jakarta parking elektronik biasanya p...
1	Bayar pakek gopaydana opo sopipay
2	Pembayaran paling gampang sebenarnya pakai Qris
3	Kalau gak di mulai tidak akan pernah mulai bis...
4	Wkwk berharap parkir motor di seluruh Ponorog...

Gambar 4. 7 Hasil Cleansing

Pada Gambar 4.7 terlihat bagian atas adalah data sebelum dibersihkan simbol yang dimuat dalam komentar kemudian di bawahnya adalah hasil pembersihan yang dimuat pada dataframe. Simbol, hastag, mention sudah tidak terlihat.

Kemudian untuk tahap *filtering* atau *stopword removal* yang berguna untuk pemilihan kata-kata yang dianggap penting. Pada tahap ini memanfaatkan kamus *stopword* dengan metode *lexicon based* dari sumber <https://raw.githubusercontent.com/masdeviD/ID-Stopwords/master/id.stopwords.02.01.2016.txt> terdiri dari 758 kata bahasa indonesia contohnya seperti berikut.

1. ada
2. adalah
3. adanya
4. adapun
5. agak
6. agaknya
7. agar
8. dst

Berikut proses yang dilakukan pada Gambar 4.8.

```
# Stopword Removal
!wget https://raw.githubusercontent.com/masdeviD/ID-Stopwords/master/id.stopwords.02.01.2016.txt
# Read stopwords file
with open('id.stopwords.02.01.2016.txt', 'r') as f:
    stopwords = f.read().splitlines()

# Function to remove stopwords
def remove_stopwords(text):
    words = text.split()
    filtered_words = [word for word in words if word.lower() not in stopwords]
    return ' '.join(filtered_words)

# Apply stopwords removal to the 'Comment' column
df['Comment'] = df['Comment'].apply(remove_stopwords)

# Display the DataFrame after stopwords removal
df.head()
```

Gambar 4. 8 Stopword Removal

Setelah melalui proses penghilangan atau *filtering* kata-kata *stopword* langkah selanjutnya dalam proses ini adalah penghilangan kata-kata slang berbahasa Indonesia menggunakan kamus yang didapatkan dari colloquial-indonesian-lexicon yang terdiri dari 3.592 kata colloquial atau dikenal dengan bahasa alay. Pada Gambar 4.9 proses menggunakan kamus yang sudah disimpan dalam directory kemudian diakses dan dipakai menggunakan fungsi `replace_slang`.

```
#Slang Remover

with open('/content/_json_colloquial-indonesian-lexicon.txt', 'r') as f:
    slang_dict = {}
    for line in f:
        # Split only on the first colon to handle potential multiple colons in the value
        if ':' in line:
            key, value = line.strip().split(':', 1)
            slang_dict[key.strip()] = value.strip()

# Fungsi untuk mengganti kata slang dengan kata baku
def replace_slang(text):
    words = text.split()
    new_words = []
    for word in words:
        if word in slang_dict:
            new_words.append(slang_dict[word])
        else:
            new_words.append(word)
    return ' '.join(new_words)

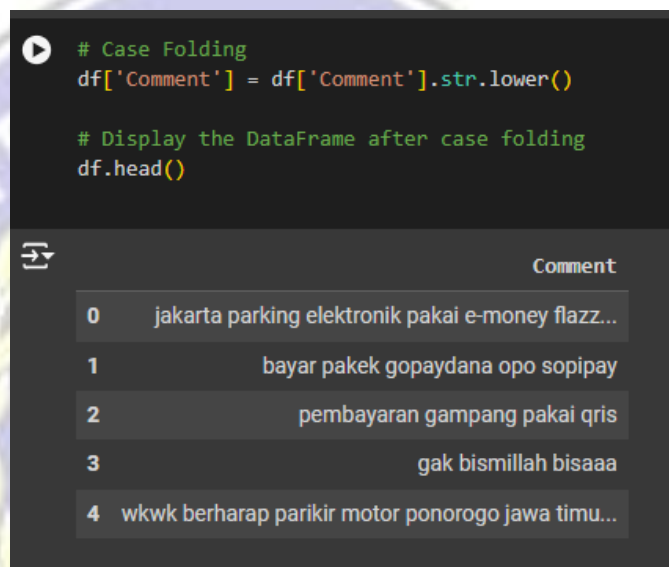
# Terapkan fungsi replace_slang ke kolom 'translated_comment'
df['Comment'] = df['Comment'].apply(replace_slang)

# Tampilkan DataFrame yang sudah diperbarui
print(df)
```

Gambar 4. 9 Slang Removal

Dataset yang dipakai pada proses-proses ini tentu sudah diubah kedalam format bahasa Indonesia. Sehingga kamus-kamus lexicon yang dipakai merupakan kumpulan kata berbahasa Indonesia.

Sebagai data yang dipakai pada proses modeling BERT perlu proses menyamaratakan huruf besar menjadi huruf kecil untuk memudahkan proses pemahaman konteks dari model menggunakan metode pada Gambar 4.10 yakni `str.lower`.



```
# Case Folding
df['Comment'] = df['Comment'].str.lower()

# Display the DataFrame after case folding
df.head()
```

	Comment
0	jakarta parking elektronik pakai e-money flazz...
1	bayar pakek gopaydana opo sopipay
2	pembayaran gampang pakai qris
3	gak bismillah bisaaa
4	wkwk berharap parkir motor ponorogo jawa timu...

Gambar 4. 10 Case Folding

Hasil dari proses pre-proccesing dari total data awal 538 menjadi 527 data bersih.

4.3. Data Labelling

Pada tahap ini peneliti memanfaatkan kamus Indonesia Sentimen terhadap dataset yang sudah diubah ke bahasa Indonesia. Skor yang didapatkan dari setiap kalimat dikategorikan menjadi 3 kelas berbeda. Berikut proses dan hasil labelisasi data ditunjukkan pada Gambar 4. 11 dan Gambar 4.12.

```
[ ] import json
import pandas as pd

# Load the positive and negative lexicons
with open('/content/_json_inset-neg.txt', 'r', encoding='utf-8') as neg_file:
    neg_lexicon = json.load(neg_file)

with open('/content/_json_inset-pos.txt', 'r', encoding='utf-8') as pos_file:
    pos_lexicon = json.load(pos_file)

# Ensure the 'comment' column exists
if 'comment' not in df.columns:
    raise ValueError("The 'comment' column is not present in the dataset.")

# Function to calculate sentiment score based on INSET Lexicon
def calculate_sentiment(text, pos_lexicon, neg_lexicon):
    words = text.lower().split()
    score = 0
    for word in words:
        score += pos_lexicon.get(word, 0) + neg_lexicon.get(word, 0)
    return score

# Function to categorize sentiment based on the score
def categorize_sentiment(score):
    if score > 0:
        return "Positive"
    elif score < 0:
        return "Negative"
    else:
        return "Neutral"

# Apply the functions to the 'comment' column
df['Sentiment_Score'] = df['comment'].apply(lambda x: calculate_sentiment(str(x), pos_lexicon, neg_lexicon))
df['Sentiment_Category'] = df['Sentiment_Score'].apply(categorize_sentiment)

# Save the labeled data
output_path = '/content/data_LABEL.csv'
df.to_csv(output_path, index=False)

print(df[['comment', 'Sentiment_Score', 'Sentiment_Category']].head())
```

Gambar 4. 11 Proses Labelisasi Data

	A	B	C	D
1	Comment	Sentiment_Score	Sentiment_Category	Sentiment_Numeric
2	kalau di jakarta parking elektronik biasanya pakai e-money flazz brizzi kartu-kartu yg l	-5 Negative		0
3	Bayar pakek gopaydana opo sopipay	-2 Negative		0
4	Pembayaran paling gampang sebenarnya pakai Qris	5 Positive		2
5	Kalau gak di mulai tidak akan pernah mulai bismillah bisaaa	0 Neutral		1
6	Wkwk berharap parkir motor di seluruh Ponorogo Jawa Timur 1000 amin	1 Positive		2
7	kadang rencana mau ngasih dua ribu trus tau ga dikembalin malih ga jadi Hehehehe	-1 Negative		0
8	Inggih kadang kita kaum ibu2 mau belanja tdk di tempat itu saja jadi klo di 5 tempat si	-8 Negative		0
9	Pakai uang tunai bisa kak bayarnya	3 Positive		2
10	Ponorogo hebat	0 Neutral		1
11	bener	2 Positive		2
12	Mantull	0 Neutral		1
13	Sing penting parkir rapi	5 Positive		2
14	oh iyo ben macak kegawe anggaran e	0 Neutral		1
15	Mbok ngnoooo ben parkir motor gak 2rb	1 Positive		2
16	Purwokerto banyumas parkir motor esih 1000	1 Positive		2
17	si paling SDM	0 Neutral		1
18	Poin pertanggung jawaban ke 3 kok sama saja gak berani tanggung jawab hanya mem	-17 Negative		0
19	jajal 10tahun kas tibake sdm e sek podo	-4 Negative		0
20	Saya suka konsep ini perlahanlahan Ponorogo menuju go digital	-5 Negative		0
21	si paling satir	-1 Negative		0
22	chuaks	0 Neutral		1
23	kasihan jukir sudah tua tidak paham dengan aplikasinya	-7 Negative		0

Gambar 4. 12 Hasil Labelisasi Data

Berdasarkan Gambar 4.12 sentimen dengan score kurang dari 0 dikategorikan negatif dan lebih dari 0 positif. Untuk memudahkan proses *training* data dikategorikan menjadi tiga 0, 1, dan 2 yang merupakan negatif, netral, positif.

4.4. Modelling

Setelah dataset diproses pada tahap *pre-processing* selanjutnya adalah menentukan parameter *tunning* untuk model sebagai berikut.

1. Epoch: 5
2. *Batch size* 16

Pemilihan *batch size* didasari karena ketersediaan sumber daya yang digunakan. *Batch size* yang lebih besar dapat mempercepat pelatihan karena lebih banyak data yang diproses secara paralel dalam satu iterasi. Namun, *batch size* yang terlalu besar juga bisa menyebabkan masalah konvergensi atau memori. Pada saat uji coba *size* yang lebih besar seperti 32 mengalami masalah ketika di jalankan.

Pada tahapan modelling ini perlu untuk melakukan inisialisasi model BERT yang sudah dilatih sebelumnya sama seperti yang dipakai untuk tokenisasi di tahap *pre-proccesing*. Model BERT yang sudah dilatih sebelumnya (*pretrained* model) digunakan untuk tugas klasifikasi urutan (*sequence classification*). Pada Gambar 4.13 memindahkan model ke perangkat yang sesuai yakni CPU, agar komputasi bisa dilakukan di sana.

```
# Initialize model
model = BertForSequenceClassification.from_pretrained('bert-base-multilingual-cased', num_labels=3)
model = model.to(device)
```

Gambar 4. 13 Inisiasi Model BERT pre-Training

Pada Gambar 4.14 upaya untuk mencegah overfitting pada model menggunakan optimizer AdamW dengan mengupdate bobot model selama pelatihan. Kemudian *learning rate* 2e-5 dan menerapkan *cross-entropy loss* untuk menghitung perbedaan antara prediksi model dan target yang diharapkan pada klasifikasi *multiclass*.

```
# Optimizer and Loss function
optimizer = AdamW(model.parameters(), lr=2e-5, correct_bias=False)
loss_fn = nn.CrossEntropyLoss().to(device)
```

Gambar 4. 14 Load Optimizer AdamW

Pelatihan dilakukan dengan beberapa skenario dengan parameter yang serupa. Pada penelitian ini data mengalami ketidakseimbangan karena jumlah data yang cenderung kecil. Peneliti kemudian melakukan upaya untuk menambah variasi dalam dataset dengan metode data augmentasi menggunakan teknik *back translation*. Data yang awalnya sejumlah 527 menjadi 1581 data dengan proses seperti pada Gambar 4.15.

```
import pandas as pd

def augment_data_with_backtranslation(df, column_name, num_augmentations):
    augmented_df = df.copy()

    for _ in range(num_augmentations):
        new_rows = []
        for _, row in df.iterrows():
            original_text = row[column_name]
            back_translated_text = back_translate(original_text)
            new_row = row.copy()
            new_row[column_name] = back_translated_text
            new_rows.append(new_row)

    # Append the newly created rows to the augmented DataFrame
    augmented_df = pd.concat([augmented_df, pd.DataFrame(new_rows)], ignore_index=True)

    return augmented_df

# Example usage:
augmented_df = augment_data_with_backtranslation(df, 'comment', 2) # Augment twice per original row
print(augmented_df)
```

	Comment	Label
0	@ricardo_kneff when parking in Jakarta usually...	Neutral
1	Payer via GoPay, Dana or ShopeePay?	Negative
2	The Easiest Payment Actually Use Qris??	Positive
3	If you don't start, it will never start. Bismi...	Negative
4	hopes to be able to park motorcycles throughou...	Positive
...	...	...
1567	It's not complicated. not yet adapted. When so...	Negative
1568	Yes, Ponorogo is the town of Reog, the area is...	Negative
1569	And all regional government offices must be mo...	Negative
1570	Let's hope paid app managers aren't confused, ...	Negative
1571	the main idea is good, but the device must be ...	Positive

[1572 rows x 2 columns]

Gambar 4. 15 Data Augmentasi

Proses yang dilakukan, dataset awal berisi 527 baris teks berbahasa indonesia pada kolom "Comment". Pada setiap iterasi augmentasi, teks asli diterjemahkan dua kali melalui proses *back translation*, sehingga menambahkan dua versi teks baru yang

teraugmentasi untuk setiap baris. Secara teknis, fungsi `augment_data_with_backtranslation` mengambil setiap baris dari DataFrame asli, menerapkan *back translation* pada teks, dan menghasilkan baris baru dengan teks yang sudah diubah. Setelah dua iterasi, jumlah baris dalam dataset meningkat menjadi 1581 (527 baris asli ditambah 1048 baris hasil augmentasi), tiga kali lipat dari ukuran awal. Hasil akhir ini kemudian disimpan sebagai DataFrame yang diperluas. Berikut rincian teknisnya.

1. Dataset Bahasa Indonesia (527 data) diterjemahkan otomatis ke Bahasa Inggris
2. Hasil Berbahasa Inggris diterjemahkan lagi ke bahasa Indonesia, hasil translasi biasanya akan berbeda dari awal namun tetap bermakna sama.
3. Setelah itu dikembalikan lagi ke bahasa inggris dan setiap baris data melakukan 2 kali iterasi sehingga berjumlah 1048 baris data hasil.
4. Hasil augmentasi tersebut kemudian dijadikan satu dengan dataset sejumlah 527 baris bahasa inggris sehingga hasil akhirnya menjadi 1581 baris data dengan variasi yang bertambah.

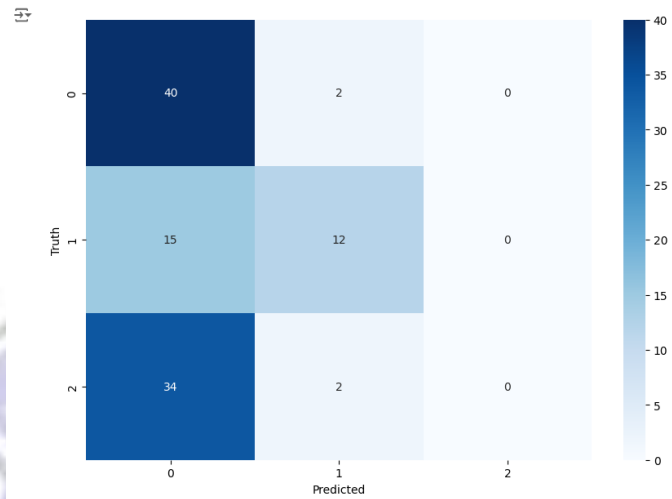
4.5. Evaluasi Model

Upaya untuk mengetahui kemampuan model dalam melakukan klasifikasi sentimen dari data komentar sosial media yang telah dikumpulkan, peneliti menggunakan *confusion matrix* pada penelitian ini. Dalam penelitian ini ada beberapa percobaan dengan parameter yang sama namun perlakuan pada data yang berbeda untuk menguji kemampuan BERT *multilingual-based* melakukan sentimen analisis pada komentar sosial media yakni dengan

4.5.1. Dataset asli(multibahasa)

Pada percobaan pertama menggunakan data asli yang mengandung bahasa jawa, Indonesia dan beberapa bahasa inggris dengan jumlah data

bersih dari proses *cleaning* kedua menjadi sebesar 527 data. Parameter yang diterapkan adalah *batch* 16 menghasilkan *accuracy* model sebagai berikut.



Gambar 4. 16 Evaluasi Model Dataset Multibahasa

Pada Gambar 4.16 memberikan gambaran hasil model dalam melakukan klasifikasi pada dataset asli menunjukkan bahwa model tidak mampu untuk mengklasifikasi kelas sentimen positif dengan label 2 pada gambar. Model dengan menggunakan BERT base *pre-trained multilingual* ternyata kurang baik dalam melakukan *text classification* pada data mayoritas berbahasa jawa.

	precision	recall	f1-score	support
0	0.45	0.95	0.61	42
1	0.75	0.44	0.56	27
2	0.00	0.00	0.00	36
accuracy			0.50	105
macro avg	0.40	0.47	0.39	105
weighted avg	0.37	0.50	0.39	105


```

[[40  2  0]
 [15 12  0]
 [34  2  0]]

```

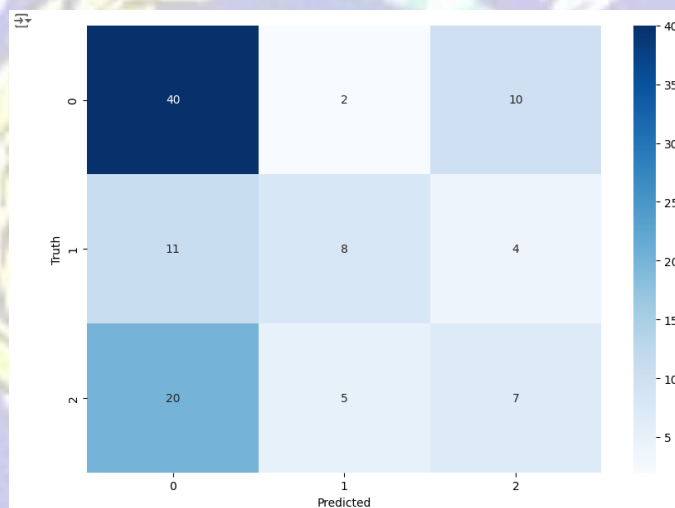
Gambar 4. 17 Hasil *Confusion matrix* Dataset Multibahasa

Dengan *Confusion matrix* model diukur bagaimana kinerjanya dalam melakukan prediksi atau klasifikasi komentar di Gambar 4.17.

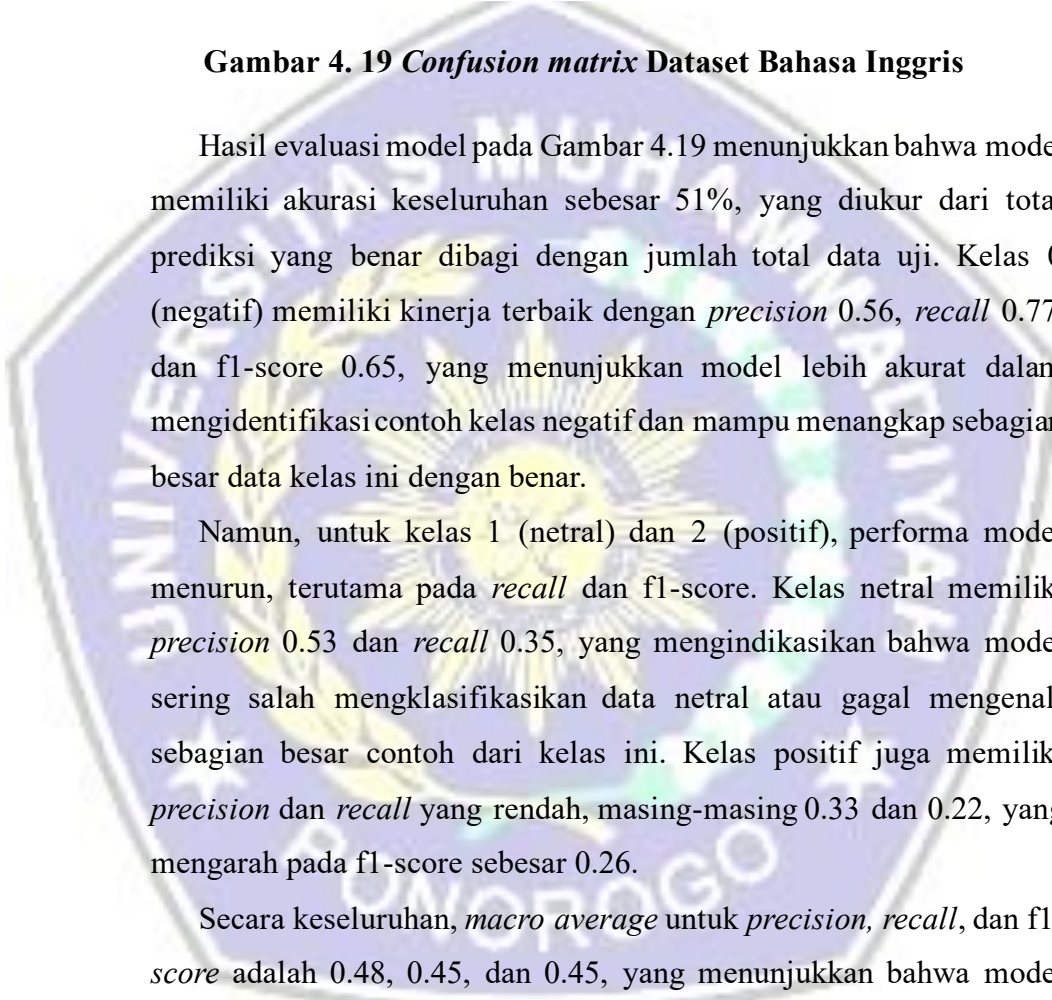
Pada data berjumlah 527 data bersih yang kemudian dibagi menjadi 80% data train dan 20% data test/validation menghasilkan akurasi model sebesar 50% menggunakan model *pre-trained multilingual* dari BERT.

4.5.2. Data diubah ke bahasa inggris

Peneliti mencoba melakukan percobaan dengan menggunakan data berbahasa inggris karena bahasa yang sangat umum pada dataset *pre-trained bert* menggunakan library *googletrans* dari python yang kemudian disempurnakan dengan pengecekan manual oleh peneliti. Pada Gambar 4.18 hasil evaluasi dengan *confusion matrix* memberikan Gambaran kemampuan klasifikasi *sentimen* oleh model tidak jauh berbeda dengan yang sebelumnya.



Gambar 4. 18 Evaluasi Model Dataset Bahasa Inggris



	precision	recall	f1-score	support
0	0.56	0.77	0.65	52
1	0.53	0.35	0.42	23
2	0.33	0.22	0.26	32
accuracy			0.51	107
macro avg	0.48	0.45	0.45	107
weighted avg	0.49	0.51	0.49	107

[[40	2	10]
[11	8	4]
[20	5	7]]

Gambar 4.19 Confusion matrix Dataset Bahasa Inggris

Hasil evaluasi model pada Gambar 4.19 menunjukkan bahwa model memiliki akurasi keseluruhan sebesar 51%, yang diukur dari total prediksi yang benar dibagi dengan jumlah total data uji. Kelas 0 (negatif) memiliki kinerja terbaik dengan *precision* 0.56, *recall* 0.77, dan *f1-score* 0.65, yang menunjukkan model lebih akurat dalam mengidentifikasi contoh kelas negatif dan mampu menangkap sebagian besar data kelas ini dengan benar.

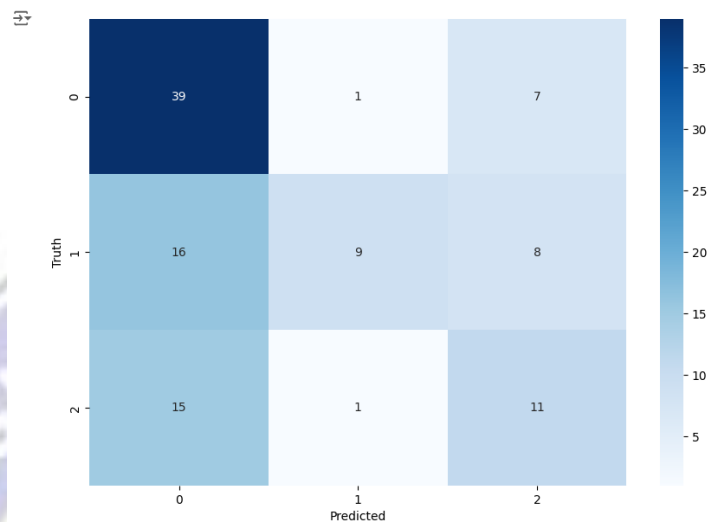
Namun, untuk kelas 1 (netral) dan 2 (positif), performa model menurun, terutama pada *recall* dan *f1-score*. Kelas netral memiliki *precision* 0.53 dan *recall* 0.35, yang mengindikasikan bahwa model sering salah mengklasifikasikan data netral atau gagal mengenali sebagian besar contoh dari kelas ini. Kelas positif juga memiliki *precision* dan *recall* yang rendah, masing-masing 0.33 dan 0.22, yang mengarah pada *f1-score* sebesar 0.26.

Secara keseluruhan, *macro average* untuk *precision*, *recall*, dan *f1-score* adalah 0.48, 0.45, dan 0.45, yang menunjukkan bahwa model memiliki kinerja yang tidak seimbang antar kelas. Peningkatan lebih lanjut dapat difokuskan pada kelas netral dan positif untuk meningkatkan *recall* dan mengurangi kesalahan klasifikasi.

4.5.3. Dataset di ubah menjadi Bahasa Indonesia

Setelah melalui percobaan sebelumnya dan hasil tidak menunjukkan akurasi model yang baik dalam data komentar sosial media yang

beragam. Peneliti melakukan percobaan ketiga Gambar 4.20 dengan mengubah bahasa dalam dataset menjadi bahasa Indonesia secara keseluruhan menggunakan *library* googlettrans dari python yang kemudian disempurnakan dengan pengecekan manual oleh peneliti.



Gambar 4. 20 Evaluasi Model Dataset Bahasa Indonesia

	precision	recall	f1-score	support
0	0.56	0.83	0.67	47
1	0.82	0.27	0.41	33
2	0.42	0.41	0.42	27
accuracy			0.55	107
macro avg	0.60	0.50	0.50	107
weighted avg	0.60	0.55	0.52	107
[[39 1 7]				
[16 9 8]				
[15 1 11]]				

Gambar 4. 21 Confusion matrix Dataset Bahasa Indonesia

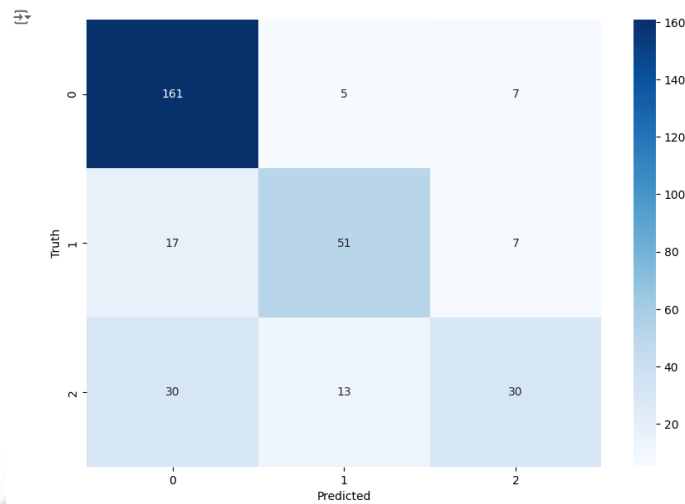
Berdasarkan hasil Gambar 4.21 akurasi keseluruhan adalah 55%, dengan performa yang bervariasi di setiap kelas. Kelas negatif (0) menunjukkan performa yang cukup baik, dengan f1-score 0.67 yang didorong oleh *recall* tinggi (0.83), meskipun *precision*-nya hanya 0.56. Kelas netral (1) memiliki *precision* tertinggi (0.82), namun rendahnya *recall* (0.27) menyebabkan f1-score hanya 0.41. Untuk kelas positif (2),

precision dan *recall* berada pada level rendah (0.42), menghasilkan f1-score yang juga rendah sebesar 0.42, menunjukkan kesulitan model dalam mengenali kelas ini. Rata-rata f1-score secara keseluruhan (macro avg) adalah 0.50, yang berarti bahwa performa model secara umum cukup seimbang di antara ketiga kelas, namun tetap ada ruang untuk perbaikan, terutama dalam hal *recall* untuk kelas netral dan positif. Selain itu, *weighted average* dari f1-score sebesar 0.52 mengindikasikan bahwa jika memperhitungkan distribusi kelas, performa model tetap tidak optimal.

Model dengan akurasi terbaik kemudian disimpan dengan metode `state_dict()` karena mempertimbangkan fleksibilitas lebih saat memuat kembali model. Metode `state_dict` hanya menyimpan bobot (parameter) model, sehingga dapat dengan mudah memuat kembali ke dalam arsitektur model yang sama atau sedikit dimodifikasi tanpa harus menyimpan seluruh objek model, yang bisa menyebabkan masalah kompatibilitas di versi PyTorch yang berbeda. Selain itu, menyimpan hanya parameter menghasilkan file yang lebih ringan dan memungkinkan untuk menggunakan kembali parameter yang telah dilatih dalam model lain atau melakukan transfer learning dengan lebih efisien.

4.5.4. Setelah data augmentasi

Pada percobaan sebelumnya dengan dataset yang diubah bahasanya namun belum menunjukkan kemampuan klasifikasi yang baik. Hal ini menunjukkan bahwa model mengalami kebingungan dalam beberapa kelas. Model memerlukan penyesuaian lebih lanjut atau tambahan data pelatihan untuk meningkatkan kinerjanya dalam mengklasifikasikan teks ke dalam kategori yang benar. Peneliti kemudian menambah variasi dataset atau data augmentasi menggunakan metode *back translation* yang menghasilkan data sebesar 1581 data. Hasil evaluasi dari model pada percobaan ini seperti Gambar 4. 22.



Gambar 4. 22 Evaluasi Model Setelah Augmentasi Data(1)

	precision	recall	f1-score	support
0	0.77	0.93	0.85	173
1	0.74	0.68	0.71	75
2	0.68	0.41	0.51	73
accuracy			0.75	321
macro avg	0.73	0.67	0.69	321
weighted avg	0.74	0.75	0.74	321
[[161 5 7]				
[17 51 7]				
[30 13 30]]				

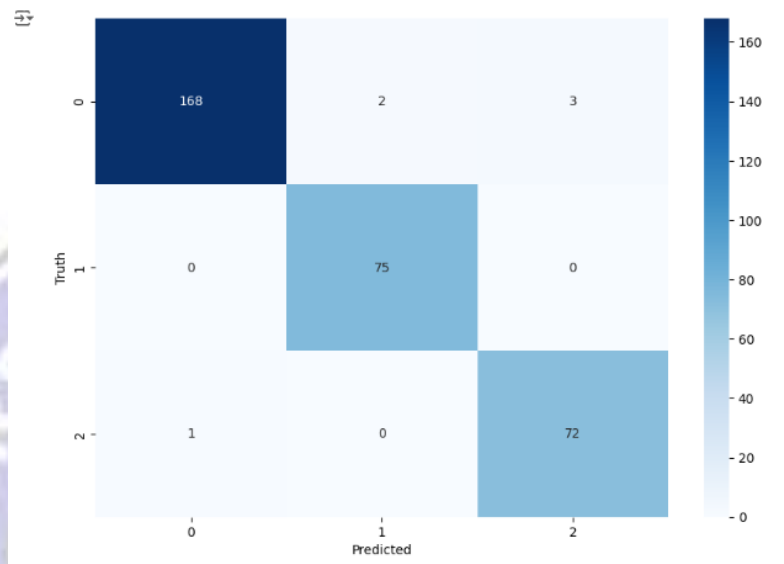
Gambar 4. 23 Confusion matrix Data Augmentasi(1)

Hasil evaluasi model dari Gambar 4.23 setelah penerapan data augmentasi menunjukkan peningkatan kinerja yang signifikan. Akurasi keseluruhan meningkat menjadi 75%, yang menunjukkan bahwa model dapat mengklasifikasikan lebih banyak data dengan benar setelah augmentasi.

Kelas negatif (0) menunjukkan performa tinggi dengan *precision* 0.77, *recall* 0.93, dan *f1-score* 0.85, mengindikasikan pengenalan yang baik dan sedikit kesalahan prediksi. Kelas netral (1) juga mengalami peningkatan dengan *f1-score* 0.71, meski *recall* masih dapat ditingkatkan. Kelas positif (2) memperlihatkan perbaikan, namun tetap rendah dengan *f1-score* 0.51. Secara keseluruhan, weighted average *f1-score*

score meningkat menjadi 0.74, menunjukkan bahwa augmentasi data memberikan dampak positif terhadap kemampuan generalisasi model.

Setelah melihat peningkatan yang signifikan dari model, peneliti selanjutnya melakukan percobaan lagi dengan data augmentasi. Berikut hasil evaluasi pada Gambar 4.24.



Gambar 4. 24 Evaluasi Model Setelah Augmentasi Data(2)

	precision	recall	f1-score	support
0	0.99	0.97	0.98	173
1	0.97	1.00	0.99	75
2	0.96	0.99	0.97	73
accuracy			0.98	321
macro avg	0.98	0.99	0.98	321
weighted avg	0.98	0.98	0.98	321
[[168 2 3]				
[0 75 0]				
[1 0 72]]				

Gambar 4. 25 Confusion matrix Data Augmentasi(2)

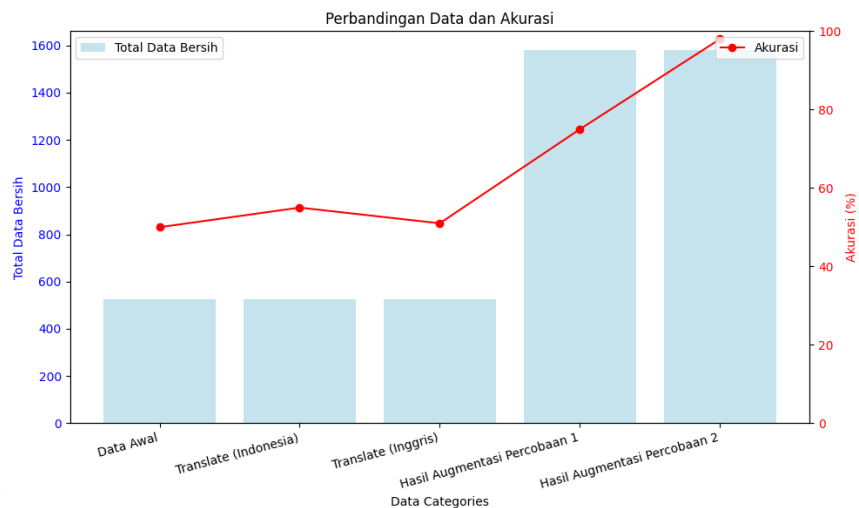
Hasil evaluasi di Gambar 4.25 *Confusion matrix* Data Augmentasi(2) menunjukkan bahwa model memiliki performa yang sangat baik dengan akurasi keseluruhan 98%. F1-score untuk masing-masing kelas juga sangat tinggi, yaitu 0.98 untuk kelas negatif(0), 0.99 untuk kelas netral (1), dan 0.97 untuk kelas positif (2), yang

menunjukkan keseimbangan antara presisi dan recall. *Confusion matrix* mengindikasikan beberapa kesalahan klasifikasi kecil, dengan dua sampel kelas negatif yang diprediksi sebagai netral dan tiga sampel kelas negatif yang diprediksi sebagai positif. Namun, secara keseluruhan, model ini sangat efisien dalam mengklasifikasikan data dengan baik.

Hasil dari seluruh percobaan dirangkum dalam Tabel 4.1 dibawah ini.

Tabel 4. 1 Hasil Evaluasi Model

Data	Bahasa	Total Data Bersih	Akurasi
<i>Data Awal</i>	<i>Jawa, Indonesia, Inggris</i>	527	50%
<i>Translate</i>	<i>Indonesia</i>	527	55%
<i>Translate</i>	<i>Inggris</i>	527	51%
<i>Hasil Augmentasi(Back translation) Percobaan 1</i>	<i>Inggris</i>	1581	75%
<i>Hasil Augmentasi(Back translation) Percobaan 2</i>	<i>Inggris</i>	1581	98%



Gambar 4. 26 Grafik modelling

Hasil dari Tabel 4.1 menunjukkan bahwa model BERT mengalami kesulitan dalam menangani data berbahasa campuran (Jawa, Indonesia, Inggris), dengan akurasi awal hanya 50%. Setelah data diterjemahkan ke dalam bahasa Indonesia dan Inggris, akurasi meningkat sedikit menjadi 55% dan 51%, menunjukkan bahwa konsistensi bahasa membantu model memahami konteks lebih baik. Namun, peningkatan signifikan terlihat setelah augmentasi data melalui *back translation*, di mana jumlah data bertambah menjadi 1581, dan akurasi melonjak hingga 98%. Ini menegaskan bahwa augmentasi data secara efektif meningkatkan performa model dalam analisis sentimen.

Berikut ini beberapa kalimat yang mengalami kesalahan klasifikasi dan yang diklasifikasi dengan benar pada Tabel 4.2 berikut ini.

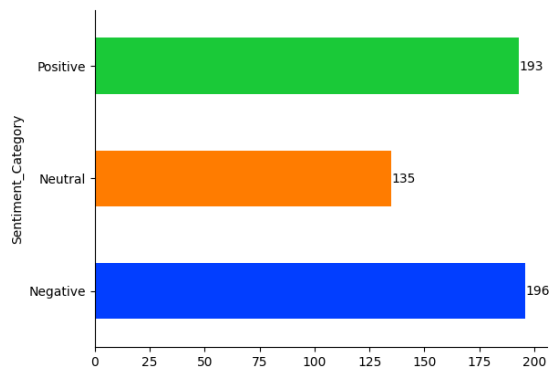
Tabel 4. 2 Klasifikasi Komentar

Komentar	Sentimen	Prediksi	Keterangan
kasihan jukir tua paham aplikasinya	Negatif	Negatif	Sesuai
baguslah jd gak merugikan masyarakat kecil& gak	Positif	Positif	Sesuai

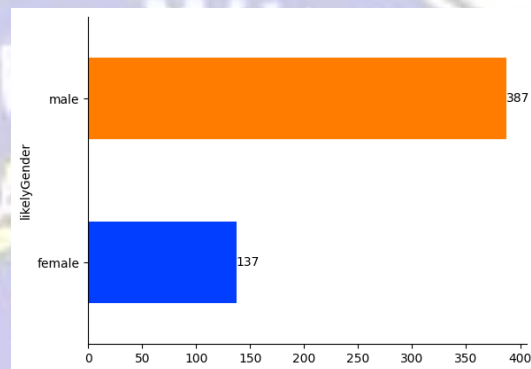
sesuka hati kasih tarif udh peraturan yg tulis jalan2 berapa2 parkir motor mobil			
bayare pake opo min tf opo pake emoney	Netral	Netral	Sesuai
payung ngarepe smp6 kae iso nganti dicem	Netral	Netral	Sesuai
suka konsep perlahanlahan ponorogo go digital	Positif	Negatif	Salah
salfok karo tulisan parkir penitipan trus nk motore ilang karcis parkire gratis	Netral	Negatif	Salah
bayar pake cash pake emoney kedepannya	Netral	Negatif	Salah
ribet eparking mending kasih fasilitas lahan khusus parkir biar gak parkir dipinggir jalan yg solutif gituloo	Negatif	Positif	Salah

4.6. Analisis dan Intepretasi

Tahapan ini dilakukan untuk mendapatkan informasi dari data komentar sosial media terkait *E-Parking* Ponorogo sebagai sumber untuk rekomendasi dan mengetahui pendapat atau opini masyarakat terkait topik terkait. Pada tahap eksplorasi data menemukan beberapa hasil yang dirangkum dalam grafik pada Gambar 4.27, Gambar 4.28 dan Gambar 4.29 berikut.

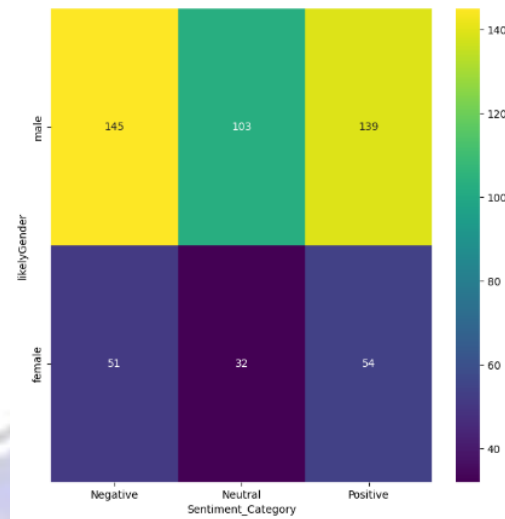


Gambar 4. 27 Distribusi Sentimen



Gambar 4. 28 Distribusi Gender

Pada Gambar 4.27 menunjukkan komentar didominasi *sentimen* negatif dan positif dari masyarakat. Kemudian pada grafik di Gambar 4.28 memberikan gambaran masyarakat yang berkomentar pada postingan terkait penerapan parkir elektronik di Kabupaten Ponorogo didominasi oleh pengguna sosial media laki-laki sebesar 387 pengguna dan pengguna sosial media Perempuan hanya berkisar 137. Dari hasil eksplorasi ini memberikan gambaran bahwa masyarakat dengan jenis kelamin laki-laki lebih memperhatikan topik ini di media sosial dibanding perempuan.



Gambar 4. 29 Distribusi Sentimen berdasarkan Gender

Untuk mengetahui lebih jelas jumlah sentimen yang dikeluarkan oleh masyarakat berdasarkan gendernya dipaparkan pada Gambar 4.29 diatas. Diagram ini menampilkan jumlah spesifik di setiap batang. Misalnya, untuk sentimen negatif, terdapat 145 laki-laki dan 51 perempuan yang mengungkapkan sentimen ini. Untuk sentimen netral, terdapat 103 laki-laki dan 32 perempuan. Untuk sentimen positif, ada 139 laki-laki dan 54 perempuan yang menunjukkan perasaan positif. Dominasi laki-laki pada setiap sentimen menunjukkan perhatian besar pada domain topik *E-Parking* yang akan diterapkan.

Interpretasi dilakukan dengan menggunakan *wordcloud* dari hasil klasifikasi sentimen. Berikut interpretasi dari hasil klasifikasi.



Gambar 4. 30 WordCloud Sentimen Negatif

Berdasarkan Gambar 4.30 menunjukkan kata-kata yang sering muncul pada komentar sosial media masyarakat dengan sentimen negatif seperti “parkir”, “jukir”, “ribet”, “terbiasa”, “rumit”, ”susah”, “sdm”, “tua” dan lainnya. Dari hasil tersebut analisis sentimen negatif terhadap penerapan parkir elektronik di kabupaten Ponorogo sebagai berikut.

1. Parkir elektonik mengalami penolakan karena dirasa ribet, rumit atau susah bagi juru parkir yang sebagian besar sudah tua dan kurang melek dengan teknologi. Serta bagi masyarakat yang belum terbiasa merasa proses parkir menjadi lebih panjang dibanding parkir manual. Beberapa komentar yang terkait seperti “mesakne jukur seng tuek pak gak paham gawe kui”, “Kesuen selak macet”, “Panggon e hospot ae buk ku bingung iki kok ndadak scan2 barang to minnn 😊”, “kesuwen, dulung i 2 ewu e gek langsung cus”,” Kesuen . . . Butue ndang di geret pak ndang di sebrangne selak okeh sing ngantri gaeyanmu akeh Noto motor harang 😊” dan lainnya.
2. Sumber daya manusia yang berada di lapangan atau juru parkir masih butuh pembiasaan atau adaptasi terkait penerapan teknologi ini dilihat dari sebagian komen yang bernarasi seperti “masih perlu adaptasi”, “Aneh aneh ae. Genah SDM e sek koyok ngene og 🤖🤖 kecepeten 10

taun”, “@tofan.ts kayane ya wis diwenehi lur, ya paling butuh waktu penyesuaian. ya muga muga isa lancar lah”, “Teknologi msa kini..... Kadang yg muda aja ada yg gaptek.. Apalgi bpk2 jukir.. Perlu ekstra utk belajar lgi... 😊”

3. Penggunaan bahu jalan atau lahan parkir yang belum memadai di lokasi uji coba juga menjadi bahasan masyarakat di sosial media terkait program parkir elektronik ini.
4. Sebagian masyarakat tidak merasa optimis atau tidak percaya dengan keberlangsungan program *E-Parking* karena beberapa uji coba program pemerintah daerah yang dibatalkan pelaksanaannya serta ketidakpercayaan masyarakat karena banyak masalah atau persoalan kabupaten yang dirasa belum dituntaskan seperti pengaturan lahan parkir, jalan desa yang tak kunjung diperbaiki dan meragukan tujuan perancangan anggaran. Beberapa komentar yang diperoleh menggambarkan analisis ini seperti berikut “Tuku alat modern bisa, jalanan berlubang nambal gak bisa. 😊”, “oh iyo ben macak kegawe anggaran e”, “halahh paling engko yo waleh, gek mbalik parkir biasa neh 😊😊😊”, “E-til ae mbalik manual,kok malah njajal liyane 😊😊”, “Teknologi melebihi akal sehat, 2023 ngurus surat2 nek kantor deso ae kon foto copy E-KTP kog parkir pake ELEKTRONIC @ponorogo.update”, dan lainnya.



Gambar 4. 31 WordCloud Sentimen Netral

Sentimen netral yang ditunjukkan pada Gambar 4.31 diatas memperlihatkan kata yang sering muncul seperti “montor”, “emoney”, ”kartu”, “opo” dan lainnya. Dari himpunan kata kata tersebut menunjukkan beberapa analisis terkait *sentimen* netral masyarakat. Didominasi pertanyaan terkait metode pembayaran akan seperti apa. Masyarakat mencoba menelaah dari postingan yang digambarkan apakah pembayaran dilakukan dengan emoney, karyu atau qr code.



Gambar 4. 32 WordCloud Sentimen Positif

Berdasarkan Gambar 4.32 menunjukkan distribusi kata pada *sentimen* positif terkait e-parkir seperti “setuju”, “penak”, “cocok”, “pungli”, “perubahan” dan lainnya.

1. Masyarakat ponorogo setuju dengan konsep ini karena memberikan transparansi tarif parkir dan jaminan parkir dibanding manual yang rawan akan pungli. Tergambarkan dari sebagian komentar seperti “Baguslah seperti ini jd kan gak merugikan masyarakat kecil... gak sesuka hati kasih tarif padahal udh ada peraturan yg di tulis di jalan2 berapa2 parkir motor dan mobil...”, “SETUJU POLL MIN, ben pak* tukang parkir i ogak sak kuarepe dw . tandang wegah bayar 2000 . jane ra sepiro sih tapi lek nyatane kendaraan e awakdewe ra diamankan ki ngwei duit parkir ki yo marai mangkel”, “Mening lki min..jelas piro2 ne tarife la MSO diweii Sewu lambene mecocong AE ” dan lainnya.
2. Konsep parkir elektronik menunjukkan progress positif Ponorogo dalam adaptasi teknologi. “@gustavianooo dijak maju gk gelem, ngko *E-Parking* wes jalan gnti komen "mbok yo ngne ket biyen" 🙏”, “Sipp mulai maju 🙏”, “Alhamdullilah ponorogo mulai maju ,, mending seperti ini soalnya saya sering di kasih uang 2000 tidak di kasih kembalian”
3. Masyarakat mendukung program ini karena merasa yang menolak dan beranggapan negatif hanya para juru parkir yang tidak mau belajar atau mencoba beradaptasi dengan sistem baru. Komentar yang diungkapkan bahkan berupa sindiran atau sarkasme terhadap juru parkir yang protes seperti “Ribet? Sg ribet kw kui wg djak penak kg angel ,,negoro kog arep maju djak pnak ae modelan ngene ki sg garai ra maju2”, “Orang yg menolak kemajuan, pasti akan mengeluarkan banyak alasan. Yg katanya ribet lah, susah lah...”, “Padine ra iso ngutel ae kok kakean alasan ribet, paling pegel ki motor ra di kapak"ne nk parkir, yo sek panggha di jaluki duwek parkir nk ngaleh” dan lain sebagainya.
4. Dukungan upaya minimalisr kebocoran PAD dari sektor parkir dengan metode parkir elektronik “Jaman digital, SDM ya harus di upgrade.

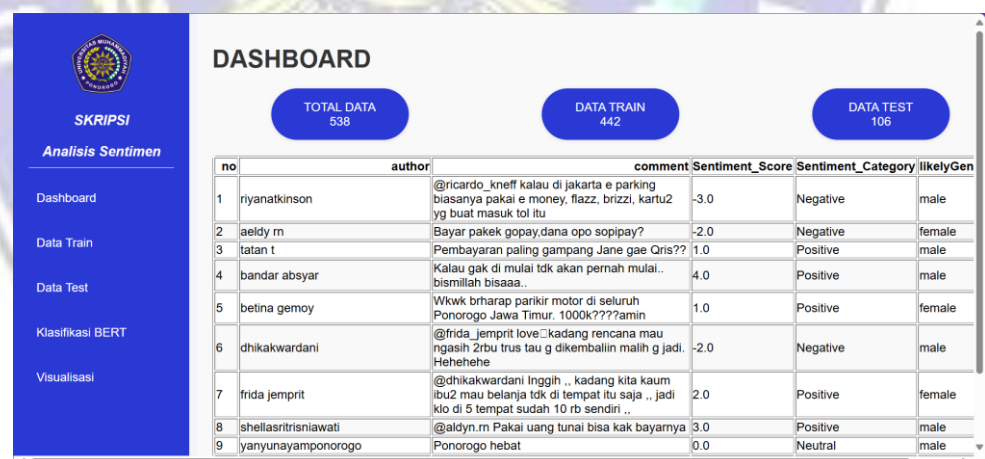
Tujuan pemkab memang modernisasi dan utk meminimalisir "kebocoran" agar PAD bertambah.”, “Pendapat ku apik lor mergo ben ora hasil dr prkir bisa langsng disetor ne nge apbd ngo bangun2 op kono ben ponorogo luweh maju dlm segala hal”, “lanjutkan e-parking,sampekl berlaku tenanan pokoke...mengurangi parkir ilegal 😄😄😄..angger mandek endi enggon kenek parkir motor 2k..basamu wes nompo dute ditinggal ngalih👉👉👉..mbok'o ditulisan motor 1k..tetep dikon mbayar 2k😄”

4.7. Graphic User Interface

Tahapan *Deployment* dalam penelitian ini dengan menampilkan hasil dari pemodelan serta analisis sentimen dengan BERT menggunakan website sederhana yang disusun oleh HTML, CSS, Python dan framework flask. Berikut *website* sederhana yang dibuat untuk memaparkan hasil sentimen analisis.

1. Halaman Dashboard

Gambar 4. 33 adalah halaman menampilkan dataset keseluruhan seperti berikut.



no	author	comment	Sentiment_Score	Sentiment_Category	likelyGen
1	riyanatkinson	@ricardo_kneff kalau di jakarta e parking biasanya pakai e money, flazz, brizzi, kartu2 yg buat masuk tol itu	-3.0	Negative	male
2	aeldy rn	Bayar pakek gopay,dana opo sopipay?	-2.0	Negative	female
3	tatan t	Pembayaran paling gampang Jane gae Qris??	1.0	Positive	male
4	bandar absyar	Kalau gak di mulai tdk akan pernah mulai.. bismillah bisaaa..	4.0	Positive	male
5	betina gemoy	Wkwk brharap parkir motor di seluruh Ponorogo Jawa Timur. 1000k???amin	1.0	Positive	female
6	dhikakwardani	@frida_jempriit love👉kadang rencana mau ngasih 2rbu trus tau g dikembalin malih g jadi. Hehehehe	-2.0	Negative	male
7	frida jempriit	@dhikakwardani Inggih .. kadang kita kaum ibu2 mau belanja tdk di tempat itu saja .. jadi klo di 5 tempat sudah 10 rb sendiri ..	2.0	Positive	female
8	shellasritrisniawati	@aldyn.rn Pakai uang tunai bisa kak bayarnya	3.0	Positive	male
9	yanyunayamponorogo	Ponorogo hebat	0.0	Neutral	male

Gambar 4. 33 Halaman Dashboard

2. Halaman data train dan test

Dataset berisi data train pada Gambar 4.34 dan data_test pada Gambar 4. 35.

DATA TRAIN	
	Comment Label
naaah benul wis jajal eparkir durung kak	0
lhadaaaah	1
iku masalahe ki nggur durung kulino ae	1
sui sui kno e toilet barang iki ngko	1
intine gelem maju opo mundur	1
pelan pelan	1
hoooh sih bener lhaaaa dadi awakmu tim manual parkir kak	0
okeeee kakak semangat	2
jukir nolak eparkir kene yo nolak mparkir podo notak e akhire gk sido mbayar	1
siplanjutkan	1
udh tua bgaimna dongksiann ga ngrti gitu	0
enak pakai e parkir kayakny	2
oke deh rpp onok kemajuan	2
mohon tukang parkir kalok niat parkir	1
yg jlasi diteribkan jukir yg ngentolan udh jlasi trtulis perda tarif parkir brp seenak e parkir sbnrnya bagus mudah udh terbiasa	2
menghindari kebocoran pad jg yuk maju yuk ponorogo hebat	1
ga berlaku	2
kadang diwei 5 ewu gak jujul alasan gak ketok ngamal	2
panggah cilik seng gae susah	0
bener bavar so dibavar karo emasan e beda wkwwk pemuunluno parkir kebanvakan aa diwei karcis karcis e maleh dadi akal	1

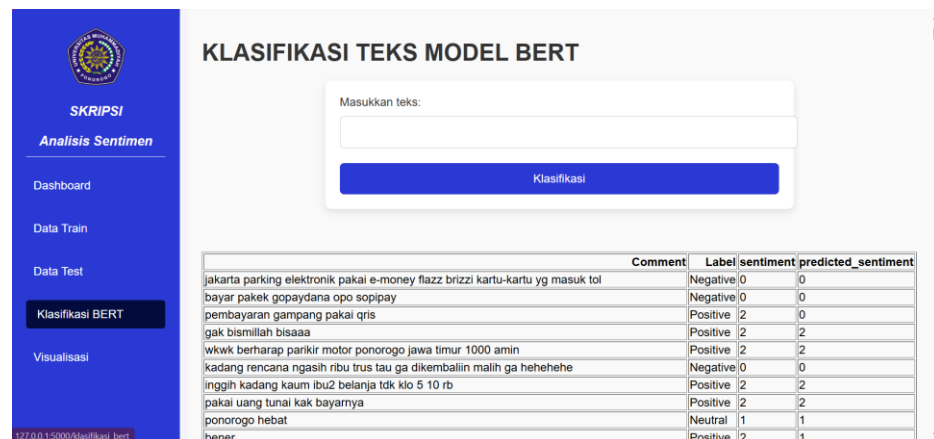
Gambar 4. 34 Halaman Data Train

DATA TEST	
	Comment Label
setuju eparkir	2
pakai ecard aja lgsung tempel gitu	2
inggh kadang kaum ibu2 belanja tdk klo 5 10 rb	2
jane lk podo bingung mending laku ae	0
jan wkwwk	1
yg lahan ndiri dibahaskn yg dbhas td ttg etilangya jls yg hubnya dg pmda wakmisalnya yo parkir ne jin hosalan2	2
namanya orang tua coba anak muda udah cepat	0
terbiasa ribet	0
meresahkan piye lur	0
orang pulang aja antri dulu nunggu beli bensin	2
mantull	1
rumit belom beradaptasi klo yg ngomong jukir merepotkan karna ngak uang gelap parkiran klo e parking sdh dipahami matilah	0
jukir nasional	2
parkir e mbayar nopo mboten	2
nahhh klo kaya gini enak jdi harga parkir ngga sndiri wkww	2
adakan pelatihan jukirnya biar sat set das des gak ngantri kaya spbu tdk pelanggan parkir nyantai yg buru2 cepat drmh 7an	2
piye kak penak sing ndi	2
pembayaran gampang pakai qris	2
team mlipir nyimak sek menowo isek enek celah	0
coboy	1

Gambar 4. 35 Halaman Data Test

3. Halaman hasil klasifikasi

Pada halaman Gambar 4.36 data hasil prediksi model ditampilkan untuk melihat kesalahan klasifikasi serta hasil evaluasi model dengan *confusion matrix* ditampilkan di halaman ini.



Gambar 4. 36 Halaman Hasil Klasifikasi

4. Halaman Visualisasi Data

Berdasarkan Gambar 4.37, data di visualkan menjadi beberapa grafik dan parameternya mulai distribusi sentimen, sebaran jumlah jenis kelamin masyarakat yang berkomentar, kemudian juga sebaran jumlah jenis sentimen komentar yang diperoleh. Selain itu *wordcloud* untuk melihat kata yang sering muncul di setiap sentimen.



Gambar 4. 37 Halaman Visualisasi Data

4.8. Hasil Pengujian Model Klasifikasi

Pada *website* dibangun *form* klasifikasi teks dengan model yang telah dibangun dengan BERT sebelumnya. Masukan teks akan menunjukkan teks tersebut masuk ke kategori sentimen positif, negatif

atau netral. Uji coba dilakukan dengan menggunakan kalimat bahasa inggris, Indonesia dan Jawa untuk melihat apakah model BERT-*multilingual* yang telah dibangun bisa akurat mengkategorikan sentimen dari teks baru yang dimasukkan. Berikut hasil uji coba yang dilakukan pada website sentimen analisis yang telah dibuat.

1. Teks Bahasa Indonesia

Teks bernarasi positif yang dimasukkan seperti Gambar 4.38 model mampu melakukan klasifikasi dengan akurat sebagai *sentimen* positif.

The screenshot shows a web application titled "KLASIFIKASI TEKS MODEL BERT". On the left is a blue sidebar with the logo of Universitas Bina Nusantara (BINUS) and a menu with items: SKRIPSI, Analisis Sentimen, Dashboard, Data Train, Data Test, Klasifikasi BERT, and Visualisasi. The main content area has a text input field with the text "saya mendukung e-parking diterapkan karena akan meningkatkan efisiensi dan tr". Below the input field is a blue "Klasifikasi" button. The results section shows "Input Text:" followed by the same input text, and "Sentimen:" followed by "Positive". At the bottom, there is a table with columns "Comment", "Label sentiment", and "predicted_sentiment". The first row contains the text "jakarta parking elektronik pakai e-money flazz brizzi kartu-kartu yg masuk tol", "Negative", and "0".

Gambar 4. 38 Teks Positif B.Indo

Kemudian untuk teks berbahasa Indonesia dengan narasi negatif juga dapat di klasifikasi secara akurat sebagai *sentimen* negatif seperti yang terlihat pada Gambar 4.39 dibawah ini.

The screenshot shows the same web application as in Gambar 4.38. The text input field now contains "kalau diterapkan di ponorogo sepertinya belum saatnya karena sdm nya masih ku". The "Klasifikasi" button is still present. The results section shows "Input Text:" followed by the input text, and "Sentimen:" followed by "Negative". The table at the bottom remains the same, with the first row showing "jakarta parking elektronik pakai e-money flazz brizzi kartu-kartu yg masuk tol", "Negative", and "0".

Gambar 4. 39 Teks Negatif B.Indo

2. Teks Bahasa Inggris

Selain uji coba dengan bahasa Indonesia selanjutnya dilakukan pengujian dengan bahasa Inggris. Pada Gambar 4.40 teks yang dimasukan adalah ungkapan positif dan berhasil di klasifikasi sebagai *sentimen* positif oleh model.

The screenshot shows a web application interface for sentiment classification. On the left is a blue sidebar with the logo of Universitas Indonesia and the text 'SKRIPSI Analisis Sentimen'. The main area has a title 'KLASIFIKASI TEKS MODEL BERT'. Below the title is a text input field with the placeholder 'Masukkan teks:'. A blue button labeled 'Klasifikasi' is below the input field. The results are displayed in a box with 'Input Text: great move, i will support this main idea' and 'Sentimen: Positive'. At the bottom, there is a table with three columns: 'Comment', 'Label/sentiment', and 'predicted_sentiment'.

Comment	Label/sentiment	predicted_sentiment
jakarta parking elektronik pakai e-money flazz brizzi kartu-kartu yg masuk tol	Negative/0	0
bayar pakek gopaydana opo sopipay	Negative/0	0

Gambar 4. 40 Teks Positif B.Ingggris

Teks yang berisi ungkapan negatif seperti Gambar 4.41 diklasifikasi dengan akurat oleh model sebagai sentimen negatif.

The screenshot shows the same web application interface as Gambar 4.40. The text input field now contains 'bad desicion i think'. The 'Klasifikasi' button is still present. The results box shows 'Input Text: bad desicion i think' and 'Sentimen: Negative'. The table at the bottom remains the same as in the previous screenshot.

Comment	Label/sentiment	predicted_sentiment
jakarta parking elektronik pakai e-money flazz brizzi kartu-kartu yg masuk tol	Negative/0	0
bayar pakek gopaydana opo sopipay	Negative/0	0

Gambar 4. 41 Teks Negatif B.Ingggris

3. Teks Bahasa Jawa

Pada penelitian ini menggunakan komentar sosial media yang banyak menggunakan bahasa jawa sehingga peneliti memutuskan untuk menguji apakah teks bahasa jawa bisa diklasifikasi oleh model. Teks bernarasi negatif pada Gambar 4.42 yang berbunyi “elek banget tak pikir-pikir” ternyata diprediksi sebagai sentimen positif oleh model yang dibangun.

The screenshot shows the 'KLASIFIKASI TEKS MODEL BERT' interface. The input text is "elek banget iki tak pikir pikir". The predicted sentiment is "Positive". Below the input area, a table displays the following data:

Comment	Label	sentiment	predicted_sentiment
jakarta parking elektronik pakai e-money flazz brizzi kartu-kartu yg masuk tol	Negative	0	0
bayar pakek gopaydana opo sopipay	Negative	0	0

Gambar 4. 42 Teks Negatif B.Jawa

Sebaliknya pada teks dengan ungkapan negatif seperti Gambar 4.43 seperti “genah marem nek ngene ki iso transparan” di klasifikasi sebagai sentimen negatif.

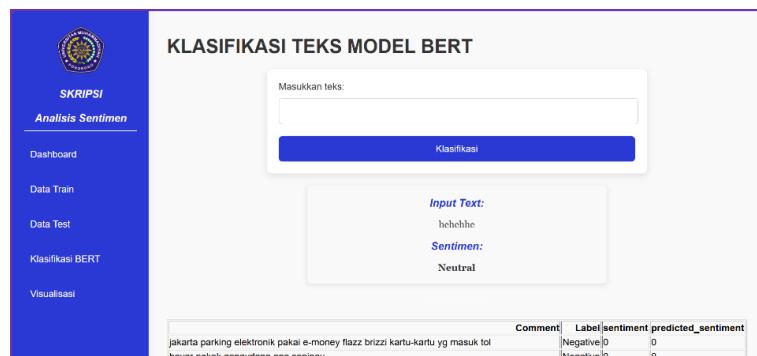
The screenshot shows the 'KLASIFIKASI TEKS MODEL BERT' interface. The input text is "genah marem nek ngene ki iso transparan". The predicted sentiment is "Negative". Below the input area, a table displays the following data:

Comment	Label	sentiment	predicted_sentiment
jakarta parking elektronik pakai e-money flazz brizzi kartu-kartu yg masuk tol	Negative	0	0
bayar pakek gopaydana opo sopipay	Negative	0	0

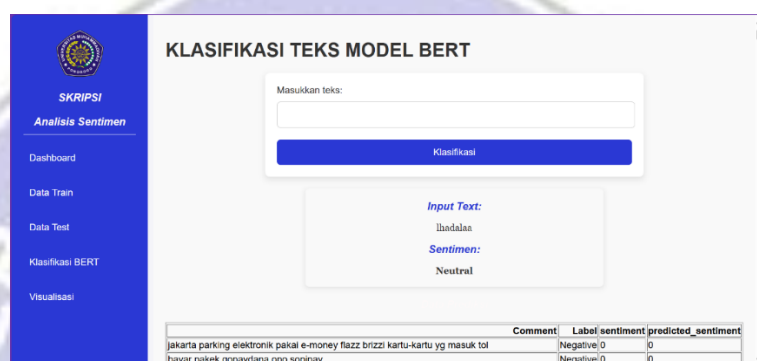
Gambar 4. 43 Teks Negatif B.Jawa

4. Teks Netral

Teks netral dimasukkan seperti Gambar 4.44 dan Gambar 4.45 Diklasifikasi dengan benar sebagai sentimen netral.



Gambar 4. 44 Teks Netral



Gambar 4. 45 Teks Netral 2

Teks yang digunakan sebagai sampel uji menggunakan pembobotan kamus Indonesia Sentimen(InSet) yang sebelumnya juga digunakan dalam tahap labelling data. Untuk teks bahasa lain mekanisme perhitungannya adalah dengan menghitung skor ketika dalam bahasa Indonesia. InSet memiliki pembobotan dari -5 hingga 5 untuk setiap kata. Setelah dihitung bobot setiap kata, teks atau kalimat pengujian kemudian ditotalkan skor akhirnya. Skor dengan nilai kurang dari 0 masuk kelas negatif, lebih dari 0 kelas positif dan 0 kelas netral.

Berikut rangkuman pada Tabel 4.3 dari tujuh sampel yang dicoba.

Tabel 4. 3 Hasil Uji Model Website

No.	Bahasa	Kalimat	Sentimen	Hasil
1.	B. Indonesia	Saya mendukung e-parking	Positif (Skor 8)	Positif

		diterapkan karena meningkatkan transparansi dan efisiensi		
2.	B.Indonesia	Kalau diterapkan di Ponorogo sepertinya belum saatnya karena SDM nya belum mumpuni	Negatif (Skor -6)	Negatif
3.	B.Ingggris	Great move, I will support this	Positif (Skor 7)	Positif
4.	B.Ingggris	Bad decision, I think	Negatif (Skor -1)	Negatif
5.	B.Jawa	EleK banget tak pikir-pikir	Negatif (Skor -6)	Positif
6.	B.Jawa	Genah marem nek ngene ki iso transparan	Positif (Skor 7)	Negatif
7.	B.Indo	hehehe	Netral (Skor 0)	Netral

Dari hasil pengujian uji klasifikasi beberapa data teks melalui website masih terjadi beberapa kesalahan seperti dalam Gambar. 4.41 dan 4.42 yang merupakan teks berbahasa jawa. Uji klasifikasi model dengan gabungan teks yang memiliki kata-kata berbahasa Indonesia dan bahasa jawa juga masih sulit membedakan jenis sentimennya. Teks yang

hanya berisi satu bahasa seperti Gambar 4.38, Gambar 4.39, Gambar 4.40 dan Gambar 4.41 mengalami keberhasilan klasifikasi. Keberhasilan uji coba klasifikasi yang dilakukan dengan beberapa teks dilihat dari hasil sentimen yang keluar apakah sesuai dengan teks yang dituliskan.

4.9. Black Box Testing

Sebagai upaya memastikan bahwa GUI yang dibuat dalam bentuk website sederhana dapat digunakan dan mampu menampilkan seluruh halaman yang dibangun dilakukan *testing* dengan *black box testing*. *Black box testing* melakukan pengujian fungsional fitur secara langsung dari hasil *website* yang dibangun. Pengujian menggunakan metrik Largest Contentful Paint (LCP) untuk mengukur kecepatan pemuatan elemen dalam website menggunakan fitur dari inspect dari browser dilakukan 3 kali kemudian dihitung rata-ratanya. Standar LCP dari browser yakni ≤ 2.5 detik Baik(cepat), 2.6 -4.0 detik perlu peningkatan, ≥ 4.0 detik buruk(lambat). Setelah diuji fungsi serta LCP dari fitur hasilnya akan diukur dengan parameter skala likert. Dengan keterangan sebagai berikut:

1. **Sangat Baik (Skor 5)** :
 - Waktu LCP ≤ 2.5 detik.
 - Fitur berfungsi 100% tanpa error.
2. **Baik (Skor 4)** :
 - Waktu LCP antara 2.6 hingga 4.0 detik.
 - Fitur bekerja dengan baik tanpa error, tetapi ada sedikit delay.
3. **Cukup (Skor 3)** :
 - Waktu LCP antara 4.1 hingga 6.0 detik.
 - Fitur berfungsi, tetapi ada keterlambatan yang cukup terasa.
4. **Buruk (Skor 2)** :
 - Waktu LCP antara 6.1 hingga 8.0 detik.

- Fitur mengalami keterlambatan yang signifikan dan kurang responsif.
5. **Sangat Buruk (Skor : 1)**
- Waktu LCP lebih dari 8 detik.
 - Fitur gagal berfungsi atau sering mengalami error.

Berikut pada Tabel 4.4 menjelaskan hasil *testing*.

Tabel 4. 4 Pengujian Black Box

No	Skenario Uji	Hasil (pass/fail)	Rata -rata LCP (detik)	Skor Likert	Keterangan
1	Navigasi ke Dashboard dari semua halaman utama (data_test, data_train, klasifikasi, visualisasi)	Pass	0.35	5	Tanpa Error Cepat
2	Navigasi ke data_train dari semua halaman utama (dashboard, data_test, klasifikasi_bert, visualisasi)	Pass	0.26	5	Tanpa Error Cepat
3	Navigasi ke data_test dari semua halaman utama (dashboard, data_train,	Pass	0.28	5	Tanpa Error Cepat

	klasifikasi_bert, visualisasi)					
4	Navigasi ke klasifikasi_bert dari semua halaman utama (dashboard, data_train, data_test, visualisasi)	Pass	0.38	5	Tanpa Error Cepat	
5	Navigasi ke visualisasi dari semua halaman utama (dashboard, data_train, data_test, klasifikasi_bert)	Pass	0.28	5	Tanpa Error Cepat	
6	Input teks pada form klasifikasi dan klik tombol Klasifikasi	Pass	1.23	5	Tanpa Error Cepat Hasil Prediksi muncul	
7	Kosongkan form klasifikasi dan klik tombol Klasifikasi	Pass	1	5	Tanpa Error Cepat Notifikasi “Please Fill out this field” muncul	
8	Input invalid teks simbol/karakter khusus (____) dan klik tombol Klasifikasi	Pass	0,31	5	Tanpa Error Cepat Hasil Prediksi muncul “Please enter a valid text.”	

BAB V

PENUTUP

5.1. Kesimpulan

Penelitian sentimen media sosial terkait E-Parking di Ponorogo menunjukkan bahwa model BERT multilingual base kurang optimal untuk data berbahasa Jawa (akurasi 50%). Penerjemahan ke Bahasa Inggris/Indonesia meningkatkan akurasi hingga 55%, sementara augmentasi data dengan back translation (1581 data) berhasil meningkatkan akurasi signifikan hingga 98%. Model dengan benar mendeteksi sentimen bahasa Inggris/Indonesia dengan baik, namun bahasa Jawa memerlukan metode tambahan seperti penerjemahan atau kamus khusus. Sentimen negatif didominasi kekhawatiran atas kerumitan e-parking, sentimen netral mencerminkan kebingungan metode pembayaran, dan sentimen positif mendukung transparansi tarif serta modernisasi teknologi.

5.2. Saran

- Disarankan untuk mempertimbangkan penggunaan model yang lebih spesifik untuk bahasa daerah seperti bahasa Jawa. Untuk metode sentimen analisis lain seperti lexicon-based yang menggunakan kamus bahasa akan cocok untuk data berbahasa daerah yang kurang biasa dengan proses modeling NLP.
- Data berbahasa daerah seperti bahasa Jawa memerlukan tindakan untuk menyesuaikan dengan kamus pra-pelatihan yang dimiliki BERT seperti melakukan penerjemahan bahasa.
- Melakukan fine-tuning pada model BERT dengan dataset yang lebih banyak dan relevan untuk meningkatkan akurasi klasifikasi sentimen.
- Penggunaan strategi augmentasi data yang lebih efektif, seperti kombinasi metode *back translation* dengan teknik augmentasi lainnya, dapat diterapkan untuk lebih meningkatkan performa model, terutama pada kelas yang memiliki kinerja rendah.

DAFTAR PUSTAKA

- [1] Y. Anjani, A. E., Purnomo, R. A., & Cahyono, "Effectiveness of Using the Parkir-Go Application from the Perspective of Market Parkirng Attendanes as an Effort to Increase Original Local Government Revenue," *J. Ekon.*, vol. 13(02), pp. 217–225, 2024, [Online]. Available: <https://ejournal.seaninstitute.or.id/index.php/Ekonomi/article/view/3976>
- [2] N. Husin, "Komparasi Algoritma Random Forest, Naïve Bayes, dan Bert Untuk Multi-Class Classification Pada Artikel Cable News Network (CNN)," *J. Esensi Infokom J. Esensi Sist. Inf. dan Sist. Komput.*, vol. 7, no. 1, pp. 75–84, 2023, doi: 10.55886/infokom.v7i1.608.
- [3] R. Fatmasari, R. K. Septiani, T. H. Pinem, D. Fabiyanto, and W. Gata, "Implementasi Algoritma BERT Pada Komentar Layanan Akademik dan Non Akademik Universitas Terbuka di Media Sosial," *Sains, Apl. Komputasi dan Teknol. Inf.*, vol. 5, no. 2, p. 96, 2024, doi: 10.30872/jsakti.v5i2.13915.
- [4] Y. Cahyono and S. Saprudin, "Analisis *Sentimen* Tweets Berbahasa Sunda Menggunakan Naive Bayes Classifier dengan Seleksi Feature Chi Squared Statistic," *J. Inform. Univ. Pamulang*, vol. 4, no. 3, p. 87, 2019, doi: 10.32493/informatika.v4i3.3186.
- [5] B. Pang and L. Lee, "Opinion Mining and *Sentimen* Analysis," *Found. Trends® Inf. Retr.*, vol. 2, no. 1–2, pp. 1–135, 2008, doi: 10.1561/15000000011.
- [6] D. Asri Y, Kuswardanu D, *MACHINE LEARNING & DEEP LEARNING: Analisis Sentimen Menggunakan Ulasan Pengguna Aplikasi*. Uwais Inspirasi Indonesia, 2019. [Online]. Available: https://books.google.co.id/books?hl=id&lr=&id=Yu7uEAAAQBAJ&oi=fnd&pg=PR9&dq=BUKU+SENTIMEN+ANALISIS&ots=mcdmdDEDgz&sig=4HP9ranwWoWtUb2LNCqTB_ik3Po&redir_esc=y#v=onepage&q=BUKU SENTIMEN ANALISIS&f=false
- [7] Ardiansyah, Adika Sri Widagdo, Krisna Nuresa Qodri, F. E. N. Saputro, and Nisrina Akbar Rizky P, "Analisis sentimen terhadap pelayanan Kesehatan berdasarkan ulasan Google Maps menggunakan BERT," *J. Fasilkom*, vol.

- 13, no. 02, pp. 326–333, 2023, doi: 10.37859/jf.v13i02.5170.
- [8] B. Kurniawan, A. Ari Aldino, and A. Rahman Isnain, “Sentimen Analisis Terhadap Kebijakan Penyelenggara Sistem Elektronik (PSE) Menggunakan Algoritma Bidirectional Encoder Representations From Transformer (BERT),” *J. Teknol. dan Sist. Inf.*, vol. 3, no. 4, pp. 98–106, 2022, [Online]. Available: <http://jim.teknokrat.ac.id/index.php/JTSI>
- [9] M. mahrus Zain, “Analisis Sentimen Pendapat Masyarakat Mengenai Vaksin Covid-19 Pada Media Sosial Twitter dengan Robustly Optimized BERT Pretraining Approach,” *J. Komput. Terap.*, vol. 7, no. 2, pp. 280–289, 2021, doi: 10.35143/jkt.v7i2.4782.
- [10] A. Jaya, “Analisis Sentimen Pandangan Public Profesi PNS (Pegawai Negeri Sipil) dari Twiter menerapkan indonesian Roberta Base Sentimen Classifier,” *Indones. J. Data Sci.*, vol. 4, no. 1, pp. 38–44, 2023, doi: 10.56705/ijodas.v4i1.66.
- [11] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning,” *Nature*, vol. 521, no. 7553, pp. 436–444, 2015, doi: 10.1038/nature14539.
- [12] A. C. Ian Goodfellow, Yoshua Bengio, *Deep learning Adaptive Computation and Machine Learning series*, Berilustra. MIT Press, 2016, 2016.
- [13] H. S. Sowmya Vajjala, Bodhisattwa Majumder, Anuj Gupta, *Practical Natural Language Processing: A Comprehensive Guide to Building Real-World NLP Systems*. O’Reilly Media, Inc, 2020. [Online]. Available: https://books.google.co.id/books?hl=id&lr=&id=hvrrDwAAQBAJ&oi=fnd&pg=PP1&dq=nlp&ots=gsev8t8Xhe&sig=aFjeanPZAkSBMeOxed_ZstvU ntk&redir_esc=y#v=onepage&q&f=false
- [14] D. Rothman, *Transformers for Natural Language Processing: Build innovative deep neural network architectures for NLP with Python, PyTorch, TensorFlow, BERT, RoBERTa, and more*. Packt Publishing Ltd., 2021.
- [15] S. Ravichandiran, *Getting Started with Google BERT: Build and train state-of-the-art natural language processing models using BERT*. Packt Publishing Ltd, 2021.
- [16] A. Vaswani *et al.*, “Attention is all you need,” *Adv. Neural Inf. Process. Syst.*,

- vol. 2017-Decem, no. Nips, pp. 5999–6009, 2017.
- [17] K. Devlin, J., Chang, M.-W., Lee, K., & Toutanova, “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding,” 2019.
- [18] L. Meyer, Y. Sun, and A. E. Martin, “Synchronous, but not entrained: exogenous and endogenous cortical rhythms of speech and language processing,” *Lang. Cogn. Neurosci.*, vol. 35, no. 9, pp. 1089–1099, 2020, doi: 10.1080/23273798.2019.1693050.
- [19] F. Koto and G. Y. Rahmaningtyas, “Inset lexicon: Evaluation of a word list for Indonesian *sentimen* analysis in microblogs,” *Proc. 2017 Int. Conf. Asian Lang. Process. IALP 2017*, vol. 2018-Janua, no. December, pp. 391–394, 2017, doi: 10.1109/IALP.2017.8300625.
- [20] R. Wirth and J. Hipp, “CRISP-DM: towards a standard process model for data mining,” *Proc. Fourth Int. Conf. Pract. Appl. Knowl. Discov. Data Min.*, no. 24959, pp. 29–39, 2000, [Online]. Available: https://www.researchgate.net/publication/239585378_CRISP-DM_Towards_a_standard_process_model_for_data_mining
- [21] G. Van Rossum and F. L. Drake, *An introduction to Python : release 2.2.2*. 2003.
- [22] M. Grinberg, *Flask Web Development*. O’Reilly Media, 2018. [Online]. Available: <https://books.google.co.id/books?id=cVIPDwAAQBAJ>
- [23] D. Setiawan, *Buku Sakti Pemrograman Web: HTML, CSS, PHP, MySQL \& Javascript*. in Anak Hebat Indonesia. Anak Hebat Indonesia, 2017. [Online]. Available: <https://books.google.co.id/books?id=HsnyDwAAQBAJ>
- [24] I. A. Shaleh, J. Prayogi, P. Pirdaus, R. Syawal, and A. Saifudin, “Pengujian Black Box pada Sistem Informasi Penjualan Buku Berbasis Web dengan Teknik Equivalent Partitions,” *Progr. Stud. Tek. Inform. Univ. Pamulang*, vol. Vol. 4 No., 2021.