

**KLASIFIKASI TANGGAPAN MASYARAKAT TERKAIT KEBIJAKAN
EFISIENSI ANGGARAN 2025 MENGGUNAKAN METODE
*LONG SHORT-TERM MEMORY***

SKRIPSI

Oleh :
HAMIDAH LUTFIYANTI MAHARANI
NIM. 210605110065



**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2025**

**KLASIFIKASI TANGGAPAN MASYARAKAT TERKAIT KEBIJAKAN
EFISIENSI ANGGARAN 2025 MENGGUNAKAN METODE
*LONG SHORT-TERM MEMORY (LSTM)***

SKRIPSI

Diajukan Kepada:
Universitas Islam Negeri Maulana Malik Ibrahim Malang
Untuk Memenuhi Salah Satu Persyaratan dalam
Memperoleh Gelar Sarjana Komputer (S.Kom)

Oleh :
HAMIDAH LUTFIYANTI MAHARANI
NIM. 210605110065

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2025**

HALAMAN PERSETUJUAN

KLASIFIKASI TANGGAPAN MASYARAKAT TERKAIT KEBIJAKAN EFISIENSI ANGGARAN 2025 MENGGUNAKAN METODE *LONG SHORT-TERM MEMORY (LSTM)*

SKRIPSI

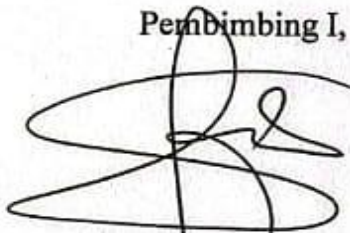
Oleh :

HAMIDAH LUTFIYANTI MAHARANI

NIM. 210605110065

Telah Diperiksa dan Disetujui untuk Diuji:
Tanggal: 17 Juni 2025

Pembimbing I,



Dr. M. Amin Hariyadi, M.T
NIP. 19670118 200501 1 001

Pembimbing II,



Dr. Zainal Abidin, M. Kom
NIP. 19760613 200501 1 004

Mengetahui,

Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi

Universitas Islam Negeri Maulana Malik Ibrahim Malang




Dr. D. Fachrul Kurniawan, M.MT., IPU
NIP. 19771020 200912 1 001

HALAMAN PENGESAHAN

KLASIFIKASI TANGGAPAN MASYARAKAT TERKAIT KEBIJAKAN EFISIENSI ANGGARAN 2025 MENGGUNAKAN METODE *LONG SHORT-TERM MEMORY (LSTM)*

SKRIPSI

Oleh :

HAMIDAH LUTFIYANTI MAHARANI
NIM. 210605110065

Telah Dipertahankan di Depan Dewan Penguji Skripsi
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan
Untuk Memperoleh Gelar Sarjana Komputer (S.Kom)
Tanggal: 23 Juni 2025

Susunan Dewan Penguji

Ketua Penguji	: <u>Dr. Irwan Budi Santoso, M.Kom</u> NIP. 19770103 201101 1 004
Anggota Penguji I	: <u>Syahiduz Zaman, M.Kom</u> NIP. 19700502 200501 1 005
Anggota Penguji II	: <u>Dr. M. Amin Hariyadi, M.T</u> NIP. 19670118 200501 1 001
Anggota Penguji III	: <u>Dr. Zainal Abidin, M.Kom</u> NIP. 19760613 200501 1 004

()
()
()
()

Mengetahui dan Mengesahkan,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang




Dr. Fachrul Kurniawan, M.MT., IPU
NIP. 19771020 200912 1 001

PERNYATAAN KEASLIAN TULISAN

Saya yang bertanda tangan di bawah ini:

Nama : Hamidah Lutfiyanti Maharani
NIM : 210605110065
Fakultas / Program Studi : Sains dan Teknologi / Teknik Informatika
Judul Skripsi : Klasifikasi Tanggapan Masyarakat Terkait Kebijakan Efisiensi Anggaran 2025 Menggunakan Metode *Long Short-Term Memory* (LSTM)

Menyatakan dengan sebenarnya bahwa Skripsi yang saya tulis ini benar-benar merupakan hasil karya saya sendiri, bukan merupakan pengambil alihan data, tulisan, atau pikiran orang lain yang saya akui sebagai hasil tulisan atau pikiran saya sendiri, kecuali dengan mencantumkan sumber cuplikan pada daftar pustaka.

Apabila dikemudian hari terbukti atau dapat dibuktikan skripsi ini merupakan hasil jiplakan, maka saya bersedia menerima sanksi atas perbuatan tersebut.

Malang, 23 Juni 2025
Yang membuat pernyataan,



Hamidah Lutfiyanti Maharani
NIM. 210605110065

MOTTO

*“Tidak Mungkin Bagi Matahari Mengejar Bulan dan Malam pun Tidak Dapat Mendahului Siang. Masing-Masing Berada Pada Garis Edarnya”
(Qs. Yasin:40)*

”Setiap Kelancaran pada langkah yang Aku Tempuh di Kehidupanku yang Sekali ini, Selalu Teriring Ridho dan Doa, Dukungan serta Semangat dari Kedua Orang Tuaku”

*”Bukan untuk Menang Kalah, tapi Tentang Bagaimana Kamu Bangkit Berkali-Kali”
| Chintya Gabriella |*

HALAMAN PERSEMBAHAN

Karya skripsi ini saya persembahkan kepada:

Zainul Arifin dan Drs. Maftukhah

Yang selalu menjadi membimbing, mengawasi serta menjadi tempat pulang yang
memberikan rasa tenang dan nyaman

Kakak-kakak saya,

Mas Zakky, Mba Nur Aminah, Mas Fani, dan Mba Yesti.

Yang telah membantu serta menjadi tempat diskusi di setiap waktu

Dosen, civitas akademik Teknik Informatika,

Yang telah memberikan arahan, bimbingan serta informasi hingga saya bisa
menamatkan studi S1 ini.

Sahabat dan Teman-Teman Seperjuangan

Teknik Informatika angkatan 2021 "Aster"

Sehat dan sukses selalu dimanapun kita berada

Serta yang Terpenting, Diri Saya Sendiri

KATA PENGANTAR

Assalamu 'alaikum warahmatullahi wabarakatuh

Dengan mengucapkan puji syukur kehadiran Allah Swt atas segala rahmat dan karunia-Nya, sehingga penulis mampu menyelesaikan penulisan skripsi ini yang berjudul “Klasifikasi Tanggapan Masyarakat Terkait Kebijakan Efisiensi Anggaran 2025 Menggunakan Metode *Long Short-Term Memory* (LSTM)” secara lancar dan tepat waktu. Penulis menyadari bahwa keberhasilan ini tidak lepas dari pertolongan-Nya. Selain itu, shawat dan salam senantiasa tercurahkan kepada Nabi Muhammad SAW, yang telah membawa umat manusia dari kegelapan menuju cahaya agama Islam yang penuh dengan rahmat, petunjuk dan kedamaian.

Skripsi ini merupakan perjalanan panjang yang penuh dengan tantangan dan pembelajaran. Proses penyelesaiannya tidak luput dari dukungan motivasi, semangat, serta bantuan moral dan material dari berbagai pihak. Untuk itu, penulis ingin menyampaikan terima kasih yang sebesar besarnya kepada:

1. Prof. Dr. M. Zainuddin, M.A., selaku rektor Universitas Islam Negeri Maulana Malik Ibrahim Malang.
2. Prof. Dr. Hj. Sri Harini, M.Si., selaku dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.
3. Dr. Ir. Fachrul Kurniawan, M.MT., IPU., selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang.
4. Dr. M. Amin Hariyadi, M.T, selaku dosen pembimbing I yang dengan penuh kesabaran dan ketulusan hati memberikan bimbingan, ilmu, memberi

semangat, arahan, masukan, serta membantu penulis dalam mengerjakan dan menyelesaikan skripsi ini.

5. Dr. Zainal Abidin, M.Kom, selaku dosen pembimbing II penulis yang selalu memberikan arahan dan masukan berharga kepada penulis dalam proses pengerjaan dan penyelesaian skripsi ini.
6. Dr. Irwan Budi Santoso, M.Kom selaku penguji I dan Syahiduz Zaman, M. Kom selaku penguji II yang telah berkenan menguji, memberikan ilmu, kritik, saran, serta memberikan masukan yang membangun sehingga penulis dapat menyelesaikan skripsi ini dengan baik.
7. Prof. Dr. Muhammad Faisal, M.T, selaku dosen yang pernah membimbing penulis dalam menyelesaikan skripsi ini dengan penuh kesabaran, keyakinan serta memberikan ilmu-Nya dan doa-doa baik untuk penulis.
8. Nia Faricha S, Si., selaku admin Program Studi Teknik Informatika yang selalu sabar membantu memberikan informasi, arahan terkait informasi akademik selama perkuliahan dan proses penulisan skripsi ini.
9. Segenap dosen, laboran, dan jajaran staff Program Studi Teknik Informatika yang telah memberikan ilmu, pengetahuan, dan dukungan selama penulis menjalani studi.
10. Orang utama yang ada dikehidupan serta menjadi panutan penulis, yaitu kedua orang tua yang tercinta, Bapak Zainul Arifin dan Ibu Drs. Maftukhah yang selalu memberikan segala usaha, doa, motivasi, dukungan, rasa aman dan nyaman, kasih sayang serta perjuangan yang tiada henti kepada penulis. Terima kasih sudah menjadi sumber semangat penulis untuk menyelesaikan

tanggung jawab ini. Tanpa doa beliau penulis tidak bisa bertahan pada kehidupan yang teramat mengejutkan ini. Penulis berharap beliau yang akan selalu ada pada setiap langkah hidup yang akan penulis jalani. Allah SWT berikan kesehatan serta perlindungan dan rezeki yang seluas luasnya.

11. Kepada kakak-kakak *support system* disegala aspek kehidupan, Ahmad Shohibuz Zakky Rosadi, Nur Aminah Kusuma Negara, Achmad Fanani Kurniawan Saputra, Rhohmini Yesti Lestari yang selalu memberikan semangat, bantuan, dukungan dan keyakinan bahwa penulis bisa menyelesaikan tantangan disetiap langkah penulis. Terima kasih sudah menjadi garda terdepan untuk selalu ada di samping penulis.
12. Teman seperjuangan "Pejuang S.Kom" yaitu Adila Qurrota A'yun dan Annisa Fitri Madani yang sedari awal perkuliahan menjadi teman, sahabat, keluarga yang memberikan berbagai bantuan serta keseruan dalam menjalani perjalanan merantau penulis. Terima kasih atas dedikasi baik dan pengalaman yang akan penulis ingat setiap momen kebersamaan-Nya
13. Teman maba penulis grup "Rwaarr" yaitu Intan Tiara Dewi, Rizqi Amalia Kartika, Dinindriya Izzatinisa, yang memberikan warna serta perjalanan indah penulis dalam menjalani setiap perkuliahan hingga selesai. Terima kasih atas kelucuan dan hiburan untuk penulis agar tetap menjalani perkuliahan dengan rasa bahagia dan senang.
14. Teman-teman penulis, Charles Iqbal, Irham DJ, Farros Shafira, Suci Wulandari, Affifah Zain, Putri Purnama, Salma Chesha, Zidan yang telah memberikan dukungan serta kekuatan kepada penulis atas penyelesaian

skripsi ini. Semoga atas apa yang sudah kita lakukan bisa menjadi ukiran kenangan indah yang akan diingat hingga nanti.

15. Sahabat penulis sejak SD hingga saat ini, alvina yang menjadi teman menugas. Serta grup "unpaedah" agis, dania dan maldha yang walaupun terpaut jarak yang jauh walaupun kita satu daerah, tetapi kekompakan masih tetap terjalin, terima kasih saling support atas apa yang sedang kita usahakan. Semoga kita semua selalu bisa berkumpul.
16. Seluruh warga Teknik Informatika khususnya angkatan 2021 "Aster" yang telah memberikan kehangatan, motivasi, semangat dan dukungan kepada penulis.
17. Kepada seluruh pihak yang terlibat, baik secara langsung maupun tidak langsung yang tidak dapat penulis sebutkan satu persatu yang kebersamaan penulis mulai dari awal perkuliahan hingga selesainya perkuliahan ini.
18. *Last but not least*, Hamidah Lutfiyanti Maharani yang telah berjuang kuat, sabar, dan berusaha sekuat tenaga dalam menyelesaikan skripsi ini. Mungkin ini awal dari kehidupan yang sesungguhnya. Tetapi terima kasih sudah bertahan dan berani ambil segala resiko yang ada. Keyakinan yang ada untuk selalu jadi lebih baik setiap harinya membuatmu sadar bahwa kamu Allah SWT ciptakan untuk menjadi makhluk yang baik di dunia dan diakhirat.

Penulis menyadari dalam penulisan serta penelitian skripsi ini tidak luput dalam kekurangan, maka dari itu penulis berterima kasih kepada semua pihak yang membantu. Penulis mengharapkan saran dan kritik yang membangun. Serta memohon maaf apabila terdapat kesalahan dalam penulis.

Malang, 20 Juni 2025

Penulis

DAFTAR ISI

HALAMAN PERSETUJUAN.....	iii
HALAMAN PENGESAHAN	iv
PERNYATAAN KEASLIAN TULISAN.....	v
MOTTO	vi
HALAMAN PERSEMBAHAN.....	vii
KATA PENGANTAR.....	viii
DAFTAR ISI	xiii
DAFTAR GAMBAR.....	xv
DAFTAR TABEL	xvi
ABSTRAK	xvii
ABSTRACT.....	xviii
مستخلص البحث.....	xix
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	7
1.3 Batasan Masalah	8
1.4 Tujuan Penelitian	8
1.5 Manfaat Penelitian	9
BAB II STUDI PUSTAKA	10
2.1 Penelitian Terkait	10
2.2 Klasifikasi Sentimen	14
2.3 <i>Preprocessing</i> Data	15
2.4 <i>Term Frequency – Inverse Document Frequency (TF-IDF)</i>	17
2.5 <i>Long Short-Term Memory (LSTM)</i>	19
BAB III METODE PENELITIAN.....	24
3.1 Rancangan Penelitian	24
3.2 Desain Sistem.....	25
3.2.1 Pengumpulan Data	26
3.2.2 Pelabelan Data	27
3.2.3 <i>Preprocessing</i> Data.....	28
3.2.4 Perhitungan TF-IDF.....	31
3.2.5 Implementasi <i>Long Short-Term Memory (LSTM)</i>	34
BAB IV HASIL DAN PEMBAHASAN	42
4.1 Skenario Pengujian.....	42
4.1.1 Data Uji	42
4.1.2 Mengukur performa sistem	44
4.2 Pemrosesan Data	46
4.3 Hasil Uji Coba dan Evaluasi	53
4.3.1 Hasil Uji Coba-1	53
4.3.2 Hasil Uji Coba-2	56
4.3.3 Hasil Uji Coba-3	58
4.4 Pembahasan.....	61
4.5 Analisis Konten.....	64
4.5.1 Tweet Kategori Positif	64
4.5.2 Tweet Kategori Negatif.....	67
4.6 Integrasi Penelitian dalam Islam	69
BAB V KESIMPULAN DAN SARAN.....	74

5.1	Kesimpulan	74
5.2	Saran	75
5.3	Rekomendasi.....	76
DAFTAR PUSTAKA.....		77

DAFTAR GAMBAR

Gambar 2.1 Struktur Jaringan Saraf Tiruan LSTM	20
Gambar 3.1 Rancangan Penelitian	24
Gambar 3.2 Skenario Pembelajaran dan Pengujian	25
Gambar 3.3 Flowchart Preprocessing Data.....	28
Gambar 3. 4 Arsitektur LSTM	35
Gambar 3. 5 Flowchart Inisialisasi Nilai.....	36
Gambar 3. 6 Flowchart Implementasi LSTM	39
Gambar 4.1 Grafik Perbandingan Data Kategori	43
Gambar 4.2 Visualisasi <i>Confusion Matrix</i>	44
Gambar 4.3 <i>Confusion Matrix Epoch 10 dan Batch Size 64 Pada Rasio 70 :30...</i>	54
Gambar 4.4 <i>Train and Validation Loss and Accuracy 70:30</i>	55
Gambar 4.5 <i>Confusion Matrix Epoch 10 dan Batch Size 32 Pada Rasio 80:20....</i>	57
Gambar 4.6 <i>Train and Validation Loss and Accuracy 80:20</i>	58
Gambar 4.7 <i>Confusion Matrix Epoch 10 dan Batch Size 32 Pada Rasio 90:10....</i>	59
Gambar 4.8 <i>Train and Validation Loss and Accuracy 90:10</i>	60

DAFTAR TABEL

Tabel 2.1 Penelitian Terkait	12
Tabel 3.1 Sampel Pelabelan Manual	27
Tabel 3.2 Contoh Proses Data Cleaning	29
Tabel 3.3 Contoh Proses Case Folding	29
Tabel 3.4 Contoh Proses Normalisasi	30
Tabel 3.5 Contoh Proses Stopwords Removal	30
Tabel 3.6 Contoh Proses Stemming	30
Tabel 3.7 Contoh Proses Tokenizing	31
Tabel 4.1 Jumlah Data Uji	42
Tabel 4.2 Pembagian Rasio Data	44
Tabel 4.3 Hasil Preprocessing <i>Data Cleaning</i>	47
Tabel 4.4 Hasil Preprocessing <i>Case Folding</i>	48
Tabel 4.5 Hasil Preprocessing Normalisasi Data	49
Tabel 4.6 Hasil Preprocessing <i>Stopword Removal</i>	50
Tabel 4.7 Hasil Preprocessing <i>Stemming</i>	51
Tabel 4.8 Hasil Preprocessing <i>Tokenizing</i>	52
Tabel 4.9 Hasil Pengujian Rasio 70:30	54
Tabel 4.10 Hasil Pelatihan Model Rasio 70:30	55
Tabel 4.11 Hasil Pengujian Rasio 80:20	56
Tabel 4.12 Hasil Pelatihan Model Rasio 80:20	57
Tabel 4.13 Hasil Pengujian Rasio 90:10	58
Tabel 4.14 Hasil Pelatihan Model Rasio 90:10	60

ABSTRAK

Maharani, Hamidah Lutfiyanti. 2025. **Klasifikasi Tanggapan Masyarakat Terkait Kebijakan Efisiensi Anggaran 2025 Menggunakan Metode *Long Short-Term Memory***. Skripsi. Program Studi Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Dr. M. Amin Hariyadi, M.T. (II) Dr. Zainal Abidin, M.Kom.

Kata Kunci: Klasifikasi Tanggapan Masyarakat, Efisiensi Anggaran 2025, TF-IDF, *Long Short-Term Memory* (LSTM).

Seiring dengan diberlakukannya kebijakan Efisiensi Anggaran 2025, pemerintah mendapatkan berbagai tanggapan dari masyarakat yang tersebar di platform media sosial. Tanggapan masyarakat ini, yang terdiri dari teks tidak terstruktur, perlu dianalisis untuk memahami sentimen publik terhadap kebijakan tersebut. Penelitian ini berfokus untuk mengklasifikasikan sentimen masyarakat, baik positif maupun negatif, menggunakan model *Long Short-Term Memory* (LSTM). Data yang digunakan berupa tweet masyarakat yang telah melalui tahapan *preprocessing*, termasuk *data cleaning*, *case folding*, normalisasi, *stopword removal*, *stemming*, dan *tokenizing*, untuk memastikan data yang digunakan relevan dan bersih dari kata-kata yang tidak perlu. Model LSTM dilatih untuk mengklasifikasikan sentimen pada tweet dan dilakukan pembagian data untuk melihat performanya. Hasilnya menunjukkan akurasi 94,38% pada rasio bagi data 80:20 dengan parameter *epoch* 10 *batch size* 32 dengan performa yang baik pada metrik presisi, *recall*, dan *f1-score*. Hasil ini membuktikan bahwa LSTM dapat digunakan secara efektif untuk menganalisis sentimen masyarakat terhadap kebijakan efisiensi anggaran 2025, memberikan insight yang bermanfaat bagi evaluasi dampak kebijakan tersebut.

ABSTRACT

Maharani, Hamidah Lutfiyanti. 2025. **Classification of Public Responses to the 2025 Budget Efficiency Policy Using the Long Short-Term Memory Method**. Thesis. Department of Informatics Engineering, Faculty of Science and Technology, State Islamic University of Maulana Malik Ibrahim Malang. Counselor: (I) Dr. M. Amin Hariyadi, M.T. (II) Dr. Zainal Abidin, M.Kom.

Along with the enactment of the Budget Efficiency 2025 policy, the government has received various responses from the public spread across social media platforms. These public responses, which consist of unstructured text, need to be analyzed to understand the public sentiment towards the policy. This research focuses on classifying public sentiments, both positive and negative, using the Long Short-Term Memory (LSTM) model, which is effective in processing sequential data such as text. The data used are public tweets that have gone through preprocessing stages, including data cleaning, case folding, normalisasi, stopword removal, stemming, and tokenizing, to ensure the data used are relevant and clean from unnecessary words. The LSTM model was trained to classify sentiment in tweets, and data splitting was performed to assess its performance. The results showed an accuracy of 94.38% on an 80:20 data ratio with parameters of 10 epochs and a batch size of 32, with good performance on the precision, recall, and F1-score metrics. These results prove that LSTM can be effectively used to analyze public sentiment towards the 2025 budget efficiency policy, providing useful insights for evaluating the impact of the policy.

Keywords: Classification of Public Responses, Budget Efficiency 2025, TF-IDF, Long Short-Term Memory (LSTM).

مستخلص البحث

مهراني، حميدة لطفيانتي. 2025. تصنيف استجابات الجمهور لسياسة كفاءة الموازنة لعام 2025 باستخدام طريقة الذاكرة طويلة وقصيرة الأجل. الأطروحة. برنامج دراسة هندسة المعلوماتية، كلية العلوم والتكنولوجيا، جامعة مولانا مالك إبراهيم الإسلامية الحكومية مالانج. المشرف: (الأول) الدكتور محمد أمين هريادي، M.T (الثاني) الدكتور زين العابدين زين العابدين، M.Kom.

الكلمات المفتاحية: تصنيف الاستجابات العامة، كفاءة الميزانية 2025، TF-IDF، الذاكرة طويلة المدى قصيرة المدى (LSTM).

مع تطبيق سياسة كفاءة الميزانية 2025، تلقت الحكومة ردود فعل متنوعة من الجمهور عبر منصات التواصل الاجتماعي. ردود فعل الجمهور هذه، التي تتكون من نصوص غير منظمة، تحتاج إلى تحليل لفهم الرأي العام تجاه هذه السياسة. تركز هذه الدراسة على تصنيف ردود فعل الجمهور، سواء كانت إيجابية أو سلبية، باستخدام نموذج Long Short-Term Memory (LSTM). البيانات المستخدمة هي تغريدات المجتمع التي خضعت لمراحل المعالجة المسبقة، بما في ذلك تنظيف البيانات، وتوحيد الأحرف، والتطبيع، وإزالة الكلمات الزائدة، واستخراج الجذور، وتجزئة الكلمات، لضمان أن البيانات المستخدمة ذات صلة وخالية من الكلمات غير الضرورية. تم تدريب نموذج LSTM على تصنيف المشاعر في التغريدات وتم تقسيم البيانات لمعرفة أدائه. أظهرت النتائج دقة بنسبة 94.38٪ على نسبة البيانات 80:20 مع معلمات epoch 10 وحجم الدفعة 32 مع أداء جيد في مقاييس الدقة والاسترجاع و f1-score. تثبت هذه النتائج أن LSTM يمكن استخدامه بشكل فعال لتحليل مشاعر المجتمع تجاه سياسة كفاءة الميزانية لعام 2025، مما يوفر رؤى مفيدة لتقييم تأثير هذه السياسة.

BAB I

PENDAHULUAN

1.1 Latar Belakang

Dalam konteks kemajuan suatu negara, pengelolaan anggaran pendapatan dan belanja menjadi elemen vital yang menentukan alokasi sumber daya dan keberhasilan pelaksanaan program-program pembangunan (Ardianti et al., 2025). Seiring dengan pentingnya pengelolaan anggaran, efisiensi dalam penggunaan dana menjadi semakin relevan untuk dilakukan agar keseimbangan fiskal tetap terjaga. Efisiensi anggaran sebagai upaya pemerintah untuk efektif dalam penggunaan dana Negara dalam menghadapi berbagai tantangan ekonomi, serta menjamin bahwa setiap rupiah yang dibelanjakan memberikan manfaat yang optimal bagi masyarakat (Choliq, 2025). Apabila anggaran digunakan secara tidak terencana dan boros, masyarakat akan merasakan dampaknya melalui hilangnya akses terhadap hak-hak dasar, seperti penurunan kualitas layanan publik.

Kebijakan Efisiensi Anggaran 2025 menandai langkah awal dalam pemerintahan Prabowo Subiyanto. Pada tanggal 22 Januari, dikeluarkan Instruksi Presiden (Inpres) Nomor 1 Tahun 2025 mengenai efisiensi belanja negara dalam pelaksanaan APBN dan APBD Tahun Anggaran 2025 (Badan Pemeriksa Keuangan, 2025). Poin pokok dari arahan Inpres tersebut mencakup pemangkasan anggaran kementerian/lembaga dan transfer ke daerah dengan penetapan target efisiensi anggaran sebesar Rp306,69 triliun. Tujuan dari pemangkasan ini adalah mengalokasikan anggaran ke dalam beberapa program

pemerintah, termasuk yang disoroti adalah Program Makan Bergizi Gratis (MBG) serta peningkatan transparansi dan akuntabilitas. Tentu saja, efisiensi anggaran ini juga dapat menimbulkan dampak besar terhadap masyarakat seperti PHK, pemotongan beasiswa dan turunya kualitas pelayanan publik (Maulana, 2025). Hal tersebut dikarenakan pemangkasan dana yang ditujukan untuk lembaga dan program sosial yang mengakibatkan berkurangnya sumber daya.

Dampak dari kebijakan tersebut menimbulkan kekhawatiran di kalangan masyarakat sebagai *stakeholder* utama mengenai kesejahteraan yang dijanjikan oleh pemerintah. Adanya ketakutan jika kebijakan ini tidak dilaksanakan dengan baik, justru akan merugikan rakyat Indonesia. Isu ini mendapatkan banyak kritik dan pendapat dari para ahli serta masyarakat, mengingat pentingnya pelaksanaan kebijakan pemerintah yang baik untuk menjamin kesejahteraan rakyat. Perkembangan teknologi, membuat masyarakat dapat menyampaikan pendapat pro dan kontranya ke berbagai media online. Kebijakan Efisiensi Anggaran 2025 ini menjadi topik yang banyak di perbincangan di platform online. Media sosial menjadi wadah diskusi terbuka bagi masyarakat mengenai kebijakan ini, dengan salah satu platform yang populer digunakan adalah X.

Sebagai salah satu platform terbesar, X menyediakan ruang bagi masyarakat untuk menyampaikan pandangan atau pendapat berbasis teks, terutama terkait kebijakan pemerintah Indonesia yang dianggap merugikan rakyat (Purwanti et al., 2024). Menurut laporan *We Are Social* dan *Meltwater*, pada April 2024, terdapat 611,3 juta pengguna X (*Twitter*) di seluruh dunia, dengan Indonesia menempati peringkat keempat. Reaksi yang muncul di platform ini mencerminkan banyak

pengguna aktif berdiskusi dan mengemukakan opini tentang kebijakan pemerintah (Mursyid & Indriyanti, 2024). X menjadi petunjuk seseorang dapat berkomunikasi, menerima, dan membagikan informasi kepada orang lain, sehingga masyarakat umum bisa mengetahuinya. Terdapat platform analisis terkait interaksi media sosial yaitu drone empiris mengatakan bahwa perbincangan efisiensi anggaran pada platform X mencapai 99,55% dari total sosial media yang ada (Rahman, 2025). Hal ini menunjukkan kebijakan efisiensi anggaran menjadi topik utama yang menarik perhatian dan melibatkan berbagai kalangan.

Pendapat masyarakat di X memberikan gambaran respon masyarakat tentang keresahan kebijakan efisiensi anggaran. Meskipun sering disampaikan secara informal dan tidak selalu terstruktur, pandangan yang muncul di X mencerminkan reaksi yang jujur, spontan, dan disampaikan secara real-time (Pravina et al., 2019). Informasi *tweet* pada X dapat digunakan untuk mengukur opini publik untuk mengetahui sentimen tentang kebijakan yang sedang terjadi saat ini. Untuk itu, penting dilakukan sebuah klasifikasi sentimen yang dapat menggambarkan pandangan dukungan atau penolakan publik mengenai kebijakan ini.

Klasifikasi sentimen membantu kita untuk memahami emosi dan pendapat yang terkandung dalam berbagai sumber informasi, seperti media sosial dan ulasan dari publik. Klasifikasi dapat dijadikan pengkategorian opini dalam teks, dengan tujuan menentukan apakah sikap penulis terhadap suatu topik bersifat positif atau negatif (Fitrana et al., 2024). Hal ini memberikan gambaran yang lebih jelas tentang reaksi masyarakat terhadap suatu isu. Dengan demikian, klasifikasi sentimen menjadi alat yang efektif untuk menganalisis persepsi publik dan dapat digunakan

sebagai evaluasi terhadap efektivitas kebijakan yang dikeluarkan pemerintah. Melalui analisis ini, pemerintah dapat mengidentifikasi kekurangan dari kebijakan tersebut dan melakukan perbaikan yang diperlukan untuk meningkatkan kesejahteraan masyarakat.

Dalam melakukan klasifikasi opini, penelitian pada perspektif Islam berkaitan dengan ayat-ayat Al-Qur'an yang membahas pembagian atau pengelompokan manusia berdasarkan ilmu dan pemahaman, seperti yang dijelaskan dalam Surah Az-Zumar ayat 9.

أَمَّنْ هُوَ قَانِتٌ آنَاءَ اللَّيْلِ سَاجِدًا وَقَائِمًا يَحْذَرُ الْآخِرَةَ وَيَرْجُوا رَحْمَةَ رَبِّهِ ۚ قُلْ هَلْ يَسْتَوِي الَّذِينَ يَعْلَمُونَ وَالَّذِينَ لَا يَعْلَمُونَ ۚ إِنَّمَا يَتَذَكَّرُ أُولَٰؤُا الْأَلْبَابِ

“(Apakah kamu hai orang musyrik yang lebih beruntung) ataukah orang yang beribadat di waktu-waktu malam dengan sujud dan berdiri, sedang ia takut kepada (azab) akhirat dan mengharapkan rahmat Tuhannya? Katakanlah: "Adakah sama orang-orang yang mengetahui dengan orang-orang yang tidak mengetahui?" Sesungguhnya orang yang berakallah yang dapat menerima pelajaran.”

(QS.Az-Zumar:9)

Berdasarkan tafsir Tahlili dalam Departemen Keagamaan Republik Indonesia (KEMENAG), ayat tersebut menjelaskan bahwa Allah SWT memerintahkan Rasul-Nya untuk bertanya kepada orang-orang kafir Mekah apakah mereka lebih beruntung daripada orang yang beribadah dengan khusyuk di malam hari, yang merasa takut akan azab Allah SWT dan berharap rahmat-Nya. Rasul diutus oleh Allah SWT untuk menanyakan apakah orang yang tahu tentang pahala dan dosa itu sama dengan orang yang tidak tahu. Orang yang tahu akan memahami akibat dari amal baik dan buruknya, sedangkan yang tidak tahu tidak memiliki harapan akan pahala atau ancaman azab. Di akhir ayat, Allah menyatakan bahwa hanya orang

yang berakal yang dapat mengambil pelajaran dari segala tanda kebesaran-Nya, baik di alam semesta maupun dalam kehidupan sehari-hari, serta dari kisah-kisah umat terdahulu.

Sesuai dengan tafsir QS. Az-Zumar ayat 9, ayat ini mengajarkan bahwa hanya orang yang berilmu dan berakal sehat yang mampu membedakan antara kebenaran dan kebatilan. Hal ini dapat diterapkan dalam konteks membedakan berbagai opini tentang efisiensi anggaran, di mana terdapat dua kubu yang mendukung (positif) dan menolak (negatif). Melalui proses klasifikasi tanggapan masyarakat, dapat berperan dalam memisahkan kebenaran dari kebatilan, sebagaimana prinsip yang diajarkan dalam ayat tersebut. Dengan demikian, hasil dari penelitian ini diharapkan dapat memberikan pemahaman yang lebih jelas mengenai sikap masyarakat terhadap kebijakan anggaran yang diterapkan pada tahun 2025.

Berbagai penelitian mengenai klasifikasi sentimen tentang kebijakan pemerintahan telah dilakukan, salah satunya oleh (Anggraini & Utami, 2021) untuk mengklasifikasikan kebijakan kartu prakerja di Indonesia dengan data *tweet* aplikasi Twitter menggunakan algoritma *Naïve Bayes*, yang menghasilkan akurasi sebesar 91.06%. Tidak hanya itu (Anggraini & Utami, 2021) juga melakukan klasifikasi terkait kebijakan vaksinasi covid-19 menggunakan algoritma *Support Vector Machine* (SVM) yang menghasilkan akurasi sebesar 89%. Dari kedua penelitian ini membuktikan bahwa melakukan klasifikasi terkait isu pemerintahan yang diambil dari data *tweet* dapat dijadikan evaluasi lebih lanjut terkait kebijakan publik.

Dalam proses klasifikasi, terdapat berbagai teknik yang digunakan untuk mencapai hasil yang optimal dalam mengidentifikasi tanggapan masyarakat. Dari

literatur sebelumnya terdapat teknik dengan metode *machine learning* konvensional yang cukup efektif yaitu *Support Vector Machine* (SVM) dan *Naïve Bayes*. Kedua metode tersebut telah terbukti mampu melakukan klasifikasi dengan baik dalam berbagai aplikasi. Tetapi kedua metode tersebut memiliki kekurangan yaitu ketergantungan terhadap *feature extraction* yang jika fitur yang relevan tidak diidentifikasi atau hilang maka performa model dapat menurun secara signifikan (Wang & Xu, 2023). Selain itu metode tersebut tidak mampu menangkap pola yang kompleks dalam data. Sehingga dibutuhkan sebuah metode untuk memperbaiki kekurangan tersebut agar mendapatkan hasil yang lebih baik lagi. Salah satu metode yang lebih canggih dan banyak digunakan dalam klasifikasi sentimen adalah *Long Short Term Memory* (LSTM).

Long Short-Term Memory (LSTM), sebuah jenis jaringan saraf tiruan yang dirancang khusus untuk menangani data sekuensial (Azrul et al., 2024). Salah satu kelebihan LSTM adalah kemampuannya mengatasi *vanishing gradient* dengan mengingat pola dan informasi penting dalam jangka panjang, yang sangat diperlukan untuk memahami konteks dalam data teks dengan alur berkelanjutan. Dengan demikian, LSTM sangat efektif dalam menganalisis teks temporal, seperti data dari aplikasi X, karena dapat mengelola informasi yang masuk melalui sel memori dan unit *gates*. Hal ini membantu LSTM untuk mempertahankan data yang relevan dalam jangka waktu lama, sementara data yang tidak relevan akan dihapus secara otomatis. Dengan keunggulan yang dimilikinya, LSTM menjadi metode yang sangat menarik untuk diterapkan dalam klasifikasi sentimen (Krisnha, 2024).

Long Short Term Memory (LSTM) dalam penentuan klasifikasi telah digunakan dalam penelitian yang dilakukan oleh (Mahmuji et al., 2023), yang membandingkan dua metode, yaitu *Naïve Bayes* dan LSTM, untuk menganalisis sentimen terkait isu penundaan Pemilu 2024. Data yang digunakan diambil dari Twitter dengan menggunakan keyword tertentu, menghasilkan 1424 data setelah dilakukan tahap *preprocessing* serta dikategorikan menjadi kelas positif dan negatif. Hasil penelitian menunjukkan bahwa metode LSTM unggul dengan akurasi mencapai 92%.

Mengacu pada permasalahan yang telah dijelaskan dengan konteks yang ada, peneliti mengajukan sebuah penelitian yaitu klasifikasi tanggapan masyarakat terkait efisiensi anggaran 2025 dengan menerapkan metode *Long Short-Term Memory* (LSTM). Tujuannya adalah untuk mengukur proporsi sentimen positif dan negatif terkait isu tersebut. Selain itu, tujuan lain dari penelitian ini adalah untuk mengkaji penerapan metode LSTM dalam mengklasifikasikan tanggapan masyarakat sekaligus menilai tingkat akurasi yang dihasilkan oleh metode tersebut.

1.2 Rumusan Masalah

Meninjau uraian latar belakang diatas, maka permasalahan rumusan masalah yang akan digunakan pada penelitian ini adalah sebagai berikut:

1. Bagaimana hasil evaluasi performa metode *Long Short-Term Memory* (LSTM) dalam mengklasifikasikan tanggapan masyarakat terkait kebijakan Efisiensi Anggaran 2025?
2. Bagaimana hasil analisis topik tanggapan masyarakat terkait kebijakan Efisiensi Anggaran 2025?

1.3 Batasan Masalah

Penelitian ini memiliki batasan-batasan yang perlu diperjelas untuk fokus pada ruang lingkup, antara lain :

1. Data yang digunakan bersumber dari platform *kaggle* dengan nama “*Budget Efficiency in Indonesia*” yang diakses pada Februari 2025.
2. Tanggapan masyarakat berupa sentimen positif dan sentimen negatif.
3. Dataset yang digunakan hanya mencakup *tweet* berupa tanggapan masyarakat terkait pembahasan “Efisiensi Anggaran 2025”, sementara *tweet* yang berbentuk informasi atau berita tidak diikutsertakan.

1.4 Tujuan Penelitian

Dalam penelitian ini, tujuan yang hendak dicapai meliputi:

1. Mengetahui hasil evaluasi performa model *Long Short-Term Memory* (LSTM) dalam mengklasifikasikan tanggapan masyarakat terkait Efisiensi Anggaran 2025.
2. Mengidentifikasi dan menganalisis topik-topik utama yang muncul dalam tanggapan masyarakat terkait kebijakan Efisiensi Anggaran 2025.

1.5 Manfaat Penelitian

Diharapkan hasil penelitian ini mampu memberikan manfaat, sebagaimana berikut:

1. Penelitian ini dapat dijadikan evaluasi pemerintah untuk memperbaiki kekurangan kebijakan Efisiensi Anggaran 2025 dengan melihat tanggapan masyarakat terkhusus di sosial media X.
2. Memperkaya literature dalam bidang klasifikasi sentiment, khususnya menggunakan model *Long Short-Term Memory* (LSTM)
3. Dapat digunakan sebagai referensi bagi penelitian selanjutnya yang ingin menerapkan metode serupa pada isu isu sosial atau kebijakan publik lainnya.

BAB II

STUDI PUSTAKA

2.1 Penelitian Terkait

Penelitian Terdahulu dengan topik atau pembahasan serupa dapat dijadikan referensi yang bisa memberikan panduan penting bagi penulis dalam menyelesaikan penelitian ini.

Penelitian yang dilakukan oleh (Deby & Isnain, 2021), membahas Analisis sentimen terhadap kebijakan *work from home* di Indonesia menggunakan *Twitter API* menunjukkan bahwa dari data yang diambil, 62.35% tweet memiliki sentimen positif, sementara 37.65% bersifat negatif. Metode yang digunakan dalam analisis ini adalah *Support Vector Machine (SVM)* dengan berbagai teknik *preprocessing*, yang bertujuan untuk meningkatkan akurasi klasifikasi sentimen. Hasil analisis menunjukkan bahwa penerapan metode ekspansi akronim, terjemahan kata slang, dan terjemahan emoji dalam tahap *preprocessing* berhasil meningkatkan nilai F1 Score hingga 83.362%.

(Jelodar et al., 2020) melakukan penelitian pada tahun 2020 yaitu mengklasifikasikan komentar terkait COVID-19 yang diambil dari platform Reddit dengan menggunakan metode gabungan dari *Word Embedding* dan *Long Short Term Memory (LSTM)*. Dengan rentang waktu yang ditentukan yaitu dari tanggal 20 Januari 2020 hingga 19 Maret 2020 didapatkan hasil data yaitu 563.769 data yang digunakan. Label sentimen ditentukan menggunakan *SentiStrength* serta divisualisasikan dengan melihat jumlah mana yang paling banyak. Tidak hanya itu

perbandingan dengan metode lain seperti SVM, *Naïve Bayes*, *Logistic Regression* dan KNN juga diterapkan. Tetapi dari hasil pengujian LSTM unggul dalam melakukan klasifikasi yaitu akurasi mencapai akurasi 81.15%, yang menunjukkan bahwa model ini efektif dalam mengklasifikasikan sentimen komentar terkait COVID-19.

Penelitian selanjutnya dilakukan oleh Purnasiwi, Kusriani dan Hanafi tahun 2023 yang menganalisis sentimen produk review skincare dengan menggunakan *Long Short Term Memory* (LSTM) dan *Word Embedding* (Purnasiwi et al., 2023). Dalam penelitian ini, peneliti menggunakan 1500 data *review* yang diambil dari website Female Daily Network. Sebelum masuk ke dalam model, peneliti membagi data dengan rasio 70:30, 80:20, dan 90:10. Selain itu, pengujian juga dilakukan untuk membandingkan kinerja LSTM dengan dan tanpa menggunakan *Word2vec*. Hasil pengujian menunjukkan akurasi sebesar 91% saat menggunakan *Word2vec*, sementara tanpa menggunakan *Word2vec* hanya mencapai 74%. Berdasarkan hasil tersebut, rasio pembagian data terbaik adalah 90:10 dengan penggunaan *Word2vec* yang menghasilkan akurasi tertinggi, yaitu 91%.

Penelitian terbaru tahun 2024, (Lisanthoni et al., 2024) melakukan penelitian yang menganalisis opini publik terhadap brand PLN. Setelah data terkumpul dari berbagai platform media digital dengan periode waktu tertentu. Untuk melabeli data ke dalam sentimen positif, negatif atau netral, peneliti menerapkan metode *Lexicon Sentiment* dengan cara menghitung polaritas berdasarkan kamus kata. Sebelum melakukan klasifikasi, data terlebih dahulu diproses melalui tahap *preprocessing* untuk mempermudah pengelolaan. Sebanyak 88.690 data berhasil dikumpulkan

dari pengambilan data ini. Kombinasi metode *Lexicon Sentiment* dan LSTM menghasilkan evaluasi model dengan menggunakan *confusion matrix*, yang menunjukkan akurasi sebesar 92%. Selain itu, nilai *precision*, *recall*, dan F1-score yang diperoleh masing-masing adalah 91,67%, 92%, dan 91,33%.

(Alfarizi et al., 2022) melakukan klasifikasi emosi yang terkandung dalam teks menggunakan metode *Long Short Term Memory* (LSTM) dan pembobotan kata TF-IDF. Peneliti melakukan percobaan untuk membandingkan metode LSTM dan LinearSVC. Berbeda dengan penelitian lainnya, penelitian ini mengklasifikasikan teks ke dalam enam kategori emosi yaitu kesedihan, kemarahan, ketakutan, cinta, kebahagiaan dan kejutan. Dataset yang digunakan berjumlah 18.000 data, yang melalui tiga tahap *preprocessing*. Setelah klasifikasi dilakukan menggunakan metode LSTM, evaluasi hasil dilakukan dengan menggunakan *confusion matrix*. Hasil penelitian menunjukkan bahwa kombinasi LSTM dan TF-IDF menghasilkan performa yang lebih baik. Secara rinci, LSTM menghasilkan akurasi 93,50%, *precision* 93%, dan *recall* 92%, sementara LinearSVC hanya menghasilkan akurasi 89%, *precision* 89%, dan *recall* 84%.

Tabel 2.1 Penelitian Terkait

No	Peneliti (tahun)	Judul Penelitian	Metode	Hasil Penelitian	Perbedaan
1.	(Deby & Isnain, 2021)	<i>Sentiment Analysis of Work from Home Activity using SVM with Randomized Search Optimization</i>	SVM, fitur <i>Preprocessing</i>	Hasil menunjukkan 62.35% dari tweet memiliki sentimen positif, sedangkan 37.65% negatif. Metode yang digunakan adalah <i>Support Vector Machine</i> (SVM), yang berhasil meningkatkan akurasi	Peneliti menggunakan model LSTM dalam menganalisis data.

No	Peneliti (tahun)	Judul Penelitian	Metode	Hasil Penelitian	Perbedaan
				klasifikasi sentimen.	
2.	(Jelodar et al., 2020)	<i>Deep Sentiment Clasification and Topic Discovery on Novel Coronavirus or COVID-19 Online Discussions:NLP Using LSTM Recurrent Neural Network Approach</i>	<i>Long Short Term Memory</i> , Fitur <i>Word Embedding</i> .	Dengan menerapkan model LSTM dan <i>word embedding</i> pada analisis dan data yang diambil dari sosial media reddit. Maka penelitian menghasilkan nilai akurasi pengujian sebesar 81,15%.	Peneliti menggunakan data <i>tweet</i> , serta dalam proses analisis menggunakan metode LSTM serta fitur TF-IDF.
3.	(Purnasiwi et al., 2023)	Analisis Sentimen Pada Review Produk Skincare Menggunakan <i>Word Embedding</i> dan Metode <i>Long Short-Term Memory</i> (LSTM)	<i>Long Short Term Memory</i> , <i>Word2vec</i>	Penelitian ini menggunakan data dari sebuah website. dengan membandingkan kinerja LSTM dengan dan tanpa menggunakan <i>Word2vec</i> . Hasil paling terbaik dari tiga rasio adalah pembagian data 90:10 dengan penggunaan <i>Word2vec</i> yang menghasilkan akurasi tertinggi, yaitu 91%.	Peneliti dalam melakukan analisis menggunakan metode LSTM serta hanya fitur TF-IDF. Data yang peniliti pakai adalah data <i>tweet</i> .
4.	(Lisanthoni et al., 2024)	Penerapan LSTM dalam Analisis Sentimen Berbasis <i>Lexicon</i> untuk Meningkatkan Sistem Pemantauan Citra PLN di Platform Digital	<i>Long Short Term Memory</i> , Berbasis <i>Lexicon</i>	Hasil penelitian menunjukkan akurasi yang baik, yaitu 92%. Data yang digunakan dari berbagai platform dan diolah menggunakan metode LSTM serta dilabeli menggunakan <i>Lexicon Sentiment</i> .	Dalam melabeli data, peneliti menggunakan pelabelan manual yang kemudian akan diverifikasi oleh validator. Untuk pengelolaannya, peneliti menerapkan metode LSTM.
5.	(Alfarizi et al., 2022)	<i>Emotional Text Classification Using TF-IDF (Term Frequency-</i>	<i>Long Short Term Memory</i> , <i>Linear SVC</i> , TF-IDF	Peneliti menggunakan enam kelas emosi. Hasil yang diperoleh	Peneliti dalam percobaan hanya menggunakan dua sentimen, yaitu sentimen positif

No	Peneliti (tahun)	Judul Penelitian	Metode	Hasil Penelitian	Perbedaan
		<i>Inverse Document Frequency) And LSTM (Long Short-Term Memory)</i>		menunjukkan bahwa model LSTM menghasilkan akurasi 93,50%, precision 93%, dan recall 92%. Berdasarkan hasil tersebut, kombinasi LSTM dan TF-IDF terbukti lebih unggul.	dan negatif, serta menerapkan metode LSTM dan TF-IDF.

Berdasarkan Tabel 2.1, terdapat berbagai penelitian tentang analisis sentimen yang dilakukan pada berbagai objek. Mengacu pada penelitian terdahulu, metode *Long Short-Term Memory* (LSTM) telah terbukti memiliki kinerja yang baik dalam pengolahan teks. Oleh karena itu, peneliti memilih metode LSTM untuk menganalisis sentimen terhadap data *tweet* mengenai Efisiensi Anggaran 2025 yang sedang ramai menjadi pembicaraan masyarakat saat ini yang belum pernah dilakukan pada penelitian terdahulu. Selain itu, peneliti juga menggunakan pembobotan kata TF-IDF mengacu pada kombinasi TF-IDF dan LSTM yang dinilai memiliki kinerja sistem yang baik. Analisis sentimen dalam penelitian ini hanya melibatkan dua kategori, yaitu sentimen positif dan sentimen negatif.

2.2 Klasifikasi Sentimen

Sentiment analysis (SA) merupakan salah satu bidang dalam *Natural Language Processing* (NLP) yang bertujuan untuk secara otomatis mengekstrak dan menganalisis sentimen serta pandangan yang terkandung dalam teks (Mao et al., 2024). Data yang diambil bisa berasal dari berbagai sumber seperti komentar di media sosial, ulasan produk, merek, layanan, isu politik atau instansi tertentu.

Dengan menganalisis komentar atau ulasan yang ditulis oleh pengguna di internet, klasifikasi sentimen membantu dalam mengidentifikasi sikap dan emosi yang mendasari tulisan tersebut.

Secara teknik, opini klasifikasi dapat dikategorikan menjadi tiga jenis yaitu positif, negatif, dan netral. Opini positif mencerminkan perasaan senang atau puas terhadap suatu produk, layanan atau kondisi, sedangkan opini negatif menunjukkan ketidakpuasan atau kekecewaan. Untuk opini netral menunjukkan kecenderungan emosional yang jelas. Dengan memahami kategori ini, klasifikasi sentiment dapat memberikan gambaran yang lebih jelas tentang bagaimana masyarakat merespons suatu isu atau produk.

Klasifikasi memiliki tujuan utama yaitu untuk memberikan wawasan mendalam bagi bisnis, pemerintah dan akademisi dalam memahami opini publik dan emosi yang terkandung dalam komentar yang beragam termasuk dalam kategori positif, negatif atau netral. Dengan kemampuan untuk menganalisis ribuan komentar secara efisien dan otomatis, hasil analisis sentimen mendukung pengambilan keputusan yang lebih tepat dan berbasis data untuk dijadikan bahan evaluasi perbaikan produk ataupun kinerjanya. Terdapat beberapa langkah yang dilakukan dalam klasifikasi sentimen yaitu pengambilan data, *text preprocessing*, pelabelan, pembobotan kata, klasifikasi dan evaluasi (Sejati et al., 2023).

2.3 *Preprocessing Data*

Preprocessing Data adalah tahap penting dalam analisis sentimen yang bertujuan mengubah data tidak terstruktur menjadi data terstruktur yang siap dianalisis (Septiawan, 2023). Data yang diperoleh biasanya masih dalam bentuk

mentah dan tidak dapat langsung digunakan. Oleh karena itu, tahap ini dilakukan untuk mempersiapkan data agar lebih berkualitas dan siap untuk dianalisis dengan baik. Terdapat serangkaian langkah yang dilakukan yaitu membersihkan data, menghilangkan *noise*, dan mengatur data sesuai dengan kebutuhan analisis (Amna et al., 2023). Berikut beberapa proses yang dilakukan dalam *preprocessing* data:

1. *Cleaning* merupakan proses penghapusan elemen elemen yang tidak diperlukan dalam teks seperti emoticon, spasi berlebih, tanda baca, angka, url atau karakter khusus yang tidak berkontribusi pada analisis.
2. *Case folding* merupakan proses mengubah semua huruf kapital dalam teks menjadi huruf kecil. Tujuannya agar memastikan konsistensi, ketika kata yang sama dapat diperlakukan sebagai entitas yang sama dalam analisis.
3. *Stopwords Removal* merupakan proses penghilangan kata yang sering muncul dalam bahasa tetapi tidak memiliki makna penting untuk analisis.
4. Normalisasi Data merupakan proses transformasi atau konversi kata-kata tidak baku (*slang words*) yang sering digunakan dalam komunikasi informal menjadi bentuk kata baku atau standar. Proses ini bertujuan untuk meningkatkan konsistensi dan akurasi dalam analisis teks, khususnya pada aplikasi seperti klasifikasi sentimen.
5. *Stemming* merupakan proses mengubah kata menjadi kata dasar. Tujuannya agar menyederhanakan teks dengan menghilangkan variasi kata, sehingga analisis dapat lebih mudah dilakukan

6. *Tokenizing* merupakan proses memecah teks menjadi unit-unit terkecil berupa kata yang disebut token agar memudahkan proses analisis elemen-elemen teks satu per satu.

2.4 *Term Frequency – Inverse Document Frequency (TF-IDF)*

Weighted word atau biasa disebut pembobotan kata adalah proses pemberian skor atau frekuensi kemunculan untuk setiap kata dalam suatu dokumen atau kumpulan dokumen (Yulita, 2021). Pembobotan kata memungkinkan untuk mengukur sejauh mana relevansi atau pentingnya suatu kata dalam konteks tertentu, serta digunakan untuk mengidentifikasi kemiripan antara kata yang satu dengan yang lainnya. Di samping itu, pembobotan kata juga diterapkan agar data dapat diproses dengan tepat oleh sebuah model.

TF-IDF merupakan salah satu metode pembobotan kata yang umum digunakan dalam pengolahan teks dan pemodelan bahasa alami. *Term Frequency – Inverse Document Frequency* adalah metode yang menggabungkan dua konsep yaitu *Term Frequency* dan *Document Frequency* (Sari et al., 2021). Sehingga hasil akhir yang didapatkan merupakan perkalian dari nilai TF dan nilai IDF. Dengan metode gabungan tersebut maka akan dilakukan dua konsep dalam pembobotan kata yaitu frekuensi kemunculan sebuah kata dalam dokumen dan frekuensi *invers* (kebalikan) dokumen yang mengandung kata tersebut.

Term Frequency (TF) atau dalam bahasa indonesia dikenal dengan frekuensi istilah, merupakan konsep pembobotan kata yang digunakan untuk menghitung seberapa sering suatu kata muncul dalam sebuah dokumen (Prasetyo et al., 2023). Asumsi dasar dari konsep ini adalah semakin sering suatu kata muncul dalam

dokumen, maka kata tersebut dianggap lebih relevan atau penting dalam konteks dokumen tersebut. Terdapat persamaan yang digunakan dalam perhitungan TF yaitu:

$$TF = \frac{\text{jumlah kemunculan kata dalam dokumen}}{\text{jumlah kata pada dokumen}} \quad (2.1)$$

Dari persamaan 2.1 perhitungan TF dilakukan dengan membagi jumlah kemunculan kata tertentu dengan total jumlah kata yang ada pada dokumen tersebut. Dengan perhitungan TF, dapat dilihat bahwa ketika sebuah kata sering muncul dalam dokumen tersebut maka bobot yang diberikan juga semakin tinggi. Namun, TF memiliki kekurangan yaitu ketidakmampuannya untuk memperhitungkan kemunculan kata dalam seluruh dokumen yang ada (Shehzad et al., 2022). Ketika sebuah kata muncul hampir disemua dokumen, kata tersebut dianggap sebagai kata yang sangat umum atau *common term* yang tidak memberikan informasi spesifik. Kata-kata seperti “dan”, “atau”, “adalah” sering muncul hampir setiap dokumen, tetapi tidak relevan dalam analisis, sehingga meskipun frekuensinya tinggi, bobot kata tersebut dianggap tidak signifikan. Untuk mengatasi kekurangan tersebut, digunakanlah konsep IDF. *Inverse Document Frequency* (IDF) atau disebut dengan Frekuensi Dokumen Terbalik, adalah konsep pembobotan kata yang hanya mengukur seberapa sering sebuah kata muncul dalam satu dokumen tanpa mempertimbangkan distribusi kata tersebut di seluruh koleksi dokumen (Prasetyo et al., 2023). Perhitungan dari IDF dapat dilihat dari persamaan berikut:

$$IDF = \log \frac{|D|}{DF} \quad (2.2)$$

Keterangan:

- | | | |
|----|-------|---------------------------------------|
| a) | IDF | : Frekuensi kemunculan kata terbalik |
| b) | D | : Jumlah total dokumen |
| c) | DF | : Jumlah dokumen yang mengandung kata |

Dalam perhitungan TF-IDF, langkah pertama adalah menghitung *Term Frequency* (TF), yaitu seberapa sering suatu kata muncul dalam dokumen. Kemudian, dihitung *Document Frequency* (DF), yaitu jumlah dokumen yang mengandung kata tersebut. Berdasarkan nilai DF, dihitung *Inverse Document Frequency* (IDF), yang memberikan bobot lebih tinggi pada kata-kata yang jarang muncul di dokumen lain. Terakhir, nilai TF-IDF dihitung dengan mengalikan nilai TF dan IDF. Hasilnya menunjukkan seberapa penting kata tersebut dalam dokumen, berdasarkan frekuensinya dan kelangkaannya dalam koleksi dokumen. Berikut merupakan persamaan perhitungan TF-IDF:

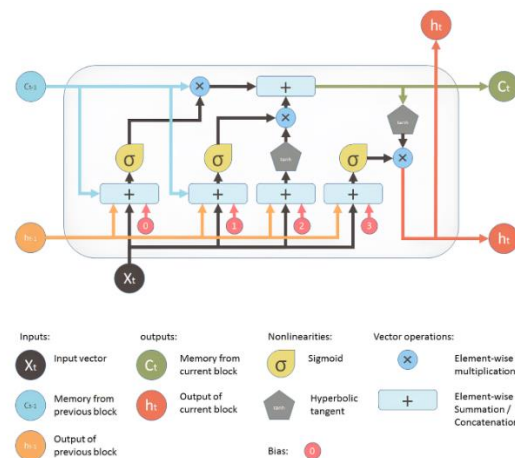
$$TF - IDF = TF \times IDF \quad (2.3)$$

2.5 Long Short-Term Memory (LSTM)

Long Short Term Memory (LSTM) pertama kali diperkenalkan oleh Sepp Hochreiter dan Jurgen Schmidhuber pada tahun 1997 untuk mengatasi masalah gradient vanishing yang sering ditemui pada RNN (*Recurrent Neural Networks*). Menurut (Nurashila et al., 2023) RNN merupakan salah satu jenis jaringan saraf tiruan (*neural network*) yang digunakan untuk mengolah informasi dari sebuah data sekuensial (seperti teks, ucapan, deret waktu). Proses dari RNN yaitu menyimpan pola informasi dari masa lalu dengan cara melakukan perulangan dalam strukturnya, sehingga informasi tersebut dapat terjaga. Namun, ketika urutan

data yang diproses semakin panjang, RNN kesulitan untuk mengingat informasi lebih lama (Wesley & Gunawan, 2024). Hal ini disebabkan oleh masalah *vanishing gradient*, dimana informasi yang datang dari langkah-langkah sebelumnya menjadi semakin kecil dan luntur seiring berjalannya waktu. Sehingga RNN tidak efektif dalam mengingat informasi jangka panjang.

Dikarenakan permasalahan tersebut, dikembangkanlah metode modifikasi dari RNN yaitu LSTM. Metode LSTM mampu menyimpan informasi penting menggunakan *memory cell* dengan waktu lebih lama dengan menggunakan cara gerbang (*gate*) untuk mengatur aliran informasi masuk dan keluar sel. Terdapat 3 jenis gate yang ada pada LSTM yaitu *input gate*, *forget gate*, *output gate* (Zahidin et al., 2024). Berikut akan dijelaskan ilustrasi struktur jaringan saraf pada LSTM



Gambar 2.1 Struktur Jaringan Saraf Tiruan LSTM
(Sumber : medium.com)

Pada gambar 2.1 merupakan ilustrasi struktur jaringan saraf tiruan dalam LSTM. LSTM memiliki struktur memori yang menyimpan informasi sepanjang urutan data yang disebut *cell*. Setiap *cell* terdiri dari dua komponen utama yaitu *cell state* dan *hidden state* (Izzadiana et al., 2023). *Cell state* berfungsi sebagai memori

jangka panjang yang mengalir melalui jaringan. Sedangkan *hidden state* merupakan output yang dihasilkan untuk digunakan dalam langkah waktu berikutnya. *Cell state* dapat diperbarui melalui transformasi linier, di mana informasi dapat ditambahkan atau dihapus menggunakan gerbang sigmoid. Gerbang pada LSTM bekerja mirip dengan jaringan saraf lainnya yang mana setiap gate juga memiliki bobot yang dipelajari selama pelatihan. Bobot ini digunakan untuk menentukan seberapa banyak informasi yang diizinkan untuk melewati gerbang.

Proses awal dalam arsitektur *Long Short Term Memory* (LSTM) dimulai dengan memilih informasi mana yang perlu disimpan atau dihapus dalam *cell state*. Bagian yang bertanggung jawab untuk memutuskan informasi yang terpilih adalah fungsi *sigmoid*. Tidak hanya itu, fungsi *sigmoid* juga menentukan bagian dari *output* sebelumnya yang harus dihapus. Fungsi *sigmoid* dijalankan di dalam *forget gate* (f_t) dengan mengambil nilai output pada waktu $t - 1$ serta nilai input waktu t . Hasil *sigmoid* yang didapatkan akan bernilai 0 dan 1. Ketika hasilnya bernilai 1, data akan disimpan dan *cell state* sebelumnya tidak akan berubah, sementara jika bernilai 0 maka *cell state* sebelumnya akan dilupakan. Menurut (Khumaidi & Nirmala, 2022) berikut merupakan persamaan dari forget gate (f_t):

$$f_t = \sigma(W_f x_t + U_f h_{t-1} + b_f) \quad (2.4)$$

Proses selanjutnya adalah memilih informasi baru yang akan ditambahkan dan disimpan ke dalam *cell state*, serta menggantikan informasi yang tidak relevan dari langkah sebelumnya melalui *input gate* (i_t). Proses ini dilakukan dalam dua tahap yaitu fungsi *sigmoid* dan fungsi *tanh* (Kurniawan et al., 2024). Pertama yaitu

fungsi *sigmoid* yang menentukan seberapa banyak informasi baru yang akan di pertahankan (1) atau di hapus (0) pada *cell state*. Kedua, fungsi *tanh* yang menghasilkan nilai kandidat dengan rentang antara -1 hingga 1 untuk ditambahkan ke dalam *cell state*. Untuk memperbarui *cell state* perhitungan dilakukan dengan mengkalikan kedua nilai fungsi. Terdapat persamaan yang digunakan sebagai berikut (Khumaidi & Nirmala, 2022):

$$i_t = \sigma(W_i x_t + U_i h_{t-1} + b_i) \quad (2.5)$$

$$\tilde{C}_t = \tanh(W_{\tilde{C}} x_t + U_{\tilde{C}} h_{t-1} + b_{\tilde{C}}) \quad (2.6)$$

Setelah informasi – informasi tersebut melewati *forget gate* dan *input gate*, maka dilakukan pembaharuan *cell state* sebelumnya (C_{t-1}) menjadi *cell state* baru C_t . Sebelum pembaharuan dilakukan, informasi yang akan dihapus pada *forget gate* dihilangkan dengan cara mengalikan *cell state* sebelumnya dengan f_t . Selanjutnya hasil perkalian tersebut akan ditambahkan dengan nilai dari *input gate* $I_t * \tilde{C}_t$ yang menghasilkan nilai baru. Dengan demikian, *cell state* baru terbentuk dan diperbarui. Jika didefinisikan ke dalam persamaan maka (Khumaidi & Nirmala, 2022) :

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t \quad (2.7)$$

Proses terakhir dalam algoritma *Long Short-Term Memory* (LSTM) adalah *output gate*, yang bertanggung jawab untuk menghasilkan output akhir dari *cell state*. Fungsi *output gate* adalah untuk menentukan bagian mana dari *cell state* yang akan dijadikan output setelah diproses. Proses dimulai dengan fungsi *sigmoid*, yang memilih bagian dari *cell state* untuk digunakan sebagai output. Kemudian, output tersebut diproses melalui lapisan *tanh* dan hasilnya dikalikan dengan output dari

fungsi *sigmoid* untuk memastikan bahwa output yang dihasilkan sesuai dengan yang telah ditentukan. Persamaan *output gate* dapat didefinisikan sebagai berikut (Khumaidi & Nirmala, 2022):

$$o_t = \sigma(W_o x_t + U_o h_{t-1} + b_o) \quad (2.8)$$

$$h_t = \tanh(C_t) \cdot o_t \quad (2.9)$$

Keterangan Persamaan:

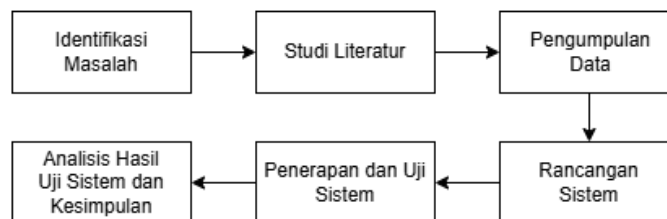
- a) f_t = *forget gate*
- b) x_t = *input data (vector input x dalam timestep t)*
- c) h_{t-1} = *vektor hidden state dalam timestep sebelumnya t-1*
- d) b_f = *vektor bias*
- e) i_t = *input gate*
- f) $\tilde{C}_t$ = *kandidat state*
- g) W_i, W_c, W_o = *matriks bobot*
- h) b_i, b_c, b_o = *vektor bias*
- i) c_t = *nilai memory cell baru*
- j) c_{t-1} = *nilai memory cell sebelumnya*
- k) o_t = *output gate*
- l) h_t = *hidden state*

BAB III

METODE PENELITIAN

3.1 Rancangan Penelitian

Rancangan penelitian yang digunakan dalam penelitian ini disusun untuk memberikan arahan yang jelas serta menetapkan target yang ingin dicapai, sehingga setiap langkah penelitian dapat dilaksanakan secara efektif. Tahapan yang akan dilakukan meliputi identifikasi masalah, rancangan sistem, hingga analisis untuk mendapatkan hasil akhir serta kesimpulan dari penelitian. Gambar 3.1 menunjukkan tahapan-tahapan perencanaan penelitian yang akan dilakukan.



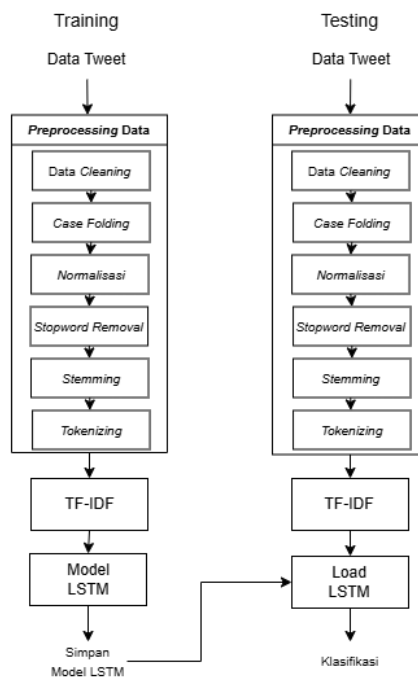
Gambar 3.1 Rancangan Penelitian

Rancangan penelitian ini diawali dengan mengidentifikasi masalah yang menjadi dasar penelitian. Untuk menjaga fokus dan kejelasan penelitian, maka ditentukan juga rumusan masalah serta batasan masalah. Setelah objek penelitian ditentukan, langkah selanjutnya adalah melakukan studi literatur dengan mencari referensi dari jurnal ilmiah atau tesis yang relevan. Dalam penelitian, data yang akan diolah perlu dikumpulkan melalui tahap pengumpulan data. Selanjutnya, dilakukan perancangan sistem sebagai alur atau syarat terkait sistem yang akan dijelaskan secara rinci agar lebih jelas dan terarah. Rancangan sistem yang telah

dibuat akan menjadi landasan dalam pengujian dan implementasi sistem. Setelah sistem dibangun, tahap pengujian dilakukan dan hasil yang diperoleh dianalisis untuk memastikan keefektifan implementasi dengan melakukan perhitungan nilai akurasi.

3.2 Desain Sistem

Sebelum melakukan penelitian, diperlukan gambaran atau perencanaan mengenai alur kerja serta interaksi antar bagian sistem. Pada rancangan sistem ini, dijelaskan tahapan mulai dari pengambilan data, pengolahan data menggunakan metode yang diterapkan dalam penelitian ini, yaitu pembobotan kata dengan TF-IDF dan *Long Short-Term Memory* (LSTM), hingga menghasilkan performa dan analisis yang diinginkan. Gambar 3.2 menggambarkan rancangan sistem dalam bentuk *flowchart* yang diterapkan dalam penelitian ini.



Gambar 3.2 Skenario Pembelajaran dan Pengujian

3.2.1 Pengumpulan Data

Dalam sebuah penelitian, dibutuhkan data atau informasi yang dikumpulkan untuk diproses dan menghasilkan analisis. Pada penelitian ini, peneliti menggunakan data sekunder yang diambil dari platform Kaggle. Dataset yang digunakan berjudul “*Budget Efficiency in Indonesia*” dengan link <https://www.kaggle.com/datasets/jocelyndumlao/budget-efficiency-in-indonesia> yang diakses pada Februari 2025. Pemilik dari Dataset ini bernama Asno Azzawagama Firdaus tetapi diupload oleh seorang pengguna *kaggle* yaitu Jocelyn Dumlao yang disediakan dan digunakan dengan izin untuk sebuah penelitian (Firdaus, 2025). Dalam dataset tersebut berisi hasil *crawling* data dari aplikasi X yang dilakukan oleh pemilik dengan memberikan informasi *tweet* terkait kebijakan tersebut. Terdapat beberapa kolom yang ada pada dataset seperti *conversation_id_str*, *created_at*, *favorite_count*, *full_text*, *id_str*, *image_url*, *in_reply_to_screen_name*, *lang*, *location*, *quote_count*, *reply_count* dan *retweet_count*. Jumlah data yang ada pada dataset sebanyak 10909 data.

Sebelum masuk ke dalam model, terdapat tahapan persiapan *dataset* yang akan dilakukan. Tahap pertama adalah seleksi kolom, di mana hanya kolom *full_text* yang digunakan karena penelitian ini berfokus pada klasifikasi tanggapan masyarakat. Tahap kedua adalah penghapusan duplikasi *tweet*, yang sering terjadi karena adanya *repost* atau *retweet* dari pengguna lain, serta hasil data yang serupa yang diperoleh melalui metode *crawling*. Tahap ketiga adalah seleksi *tweet*, yang bertujuan untuk menyaring *tweet* yang tidak mengandung tanggapan masyarakat atau tidak secara langsung membahas “Efisiensi Anggaran 2025”. *Tweet* yang

berisi informasi atau berita akan dibersihkan untuk mempermudah proses pelabelan dan menjaga fokus penelitian. Dari hasil tersebut didapatkan 2222 dataset tweet terkait Efisiensi Anggaran 2025. Setelah data berhasil dikumpulkan, data disimpan dalam bentuk .csv untuk mempermudah pengolahan lebih lanjut.

3.2.2 Pelabelan Data

Setelah data diperoleh, tahap selanjutnya adalah pelabelan data. Tujuan dari tahap ini adalah untuk memberikan label pada setiap data tweet dengan dua kategori yaitu positif dan negatif. Pelabelan dilakukan secara manual oleh peneliti dengan mengamati tanggapan yang terkandung dalam setiap *tweet*. Setelah label diberikan, kebenaran pelabelan akan divalidasi oleh seorang ahli bahasa yang memiliki pemahaman mendalam tentang penggunaan bahasa yang baik dan benar. Hal ini dilakukan untuk memastikan bahwa tweet tersebut benar-benar termasuk dalam kategori sentimen positif atau negatif. Pada penelitian ini, ahli bahasa yang melakukan validasi pelabelan adalah Dra. Hj. Maftukhah, yang berprofesi sebagai guru Bahasa Indonesia di SMA Islam Kepanjen. Data yang telah diberi label ini akan digunakan sebagai *dataset* untuk pengujian sistem, yang kemudian dibagi menjadi dua bagian yaitu data *training* dan data *testing*. Tabel 3.1 merupakan contoh *tweets* yang sudah dilakukan pelabelan manual dan siap untuk dilakukan pada tahapan selanjutnya.

Tabel 3.1 Sampel Pelabelan Manual

<i>Tweets</i>	Pelabelan Manual
drpd perpustakaan tutup hari minggu krn efisiensi mending hapus aja jabatan wapres biar ga ada anggaran yg kebuang untuk susu greenfield	Negatif
Efisiensi anggaran mendukung berbagai program sosial dan pembangunan nasional #AnggaranBijakUntukRakyat	Positif

@Mdy_Asmara1701 Pemangkasan anggaran itu memang utk efisiensi dan penghematan. Program Makan Bergizi Gratis itu sdh ada anggaran tersendiri. Plis mari lbh bijak menilai.	Positif
Ya Allah apakah efisiensi anggaran ini mempengaruhi toko oren juga?? Voucher diskonnya pelit bgtttttt minimal beli 100k.	Negatif
@raffi1133 Efisiensi anggaran layanan publik makin berkualitas!	Positif
DAMPAK EFISIENSI ANGGARAN INI ! @prabowo Dampaknya ke Rakyat yang butuh cari biaya hidup. Ttp yg mengherankan justru terus memperbanyak mengangkat STAFSUS yang TIDAK ADA URGENSINYA https://t.co/Y2OjtNKH9N	Negatif

3.2.3 Preprocessing Data

Preprocessing data adalah tahapan yang dilakukan untuk mengubah data mentah menjadi data yang siap diolah oleh model, guna mengetahui klasifikasi atau performa model tersebut. Dengan melakukan *preprocessing*, hasil klasifikasi dan performa model akan lebih baik, karena data yang telah diproses akan lebih bersih dan terstruktur dibandingkan dengan data mentah yang belum diproses. Pada penelitian ini *preprocessing* data dilakukan dengan bahasa pemrograman *Python*. Berikut adalah beberapa tahapan yang dilakukan pada proses *preprocessing* data.



Gambar 3.3 Flowchart Preprocessing Data

3.2.3.1 Data Cleaning

Data cleaning adalah proses membersihkan data dengan menghapus bagian-bagian yang tidak relevan atau tidak berguna, seperti karakter-karakter yang tidak standar, angka, tanda baca, emotikon, simbol, spasi berlebih, tautan, atau tagar yang

tidak memiliki kaitan dengan informasi yang ingin dianalisis. Proses ini bertujuan untuk menghasilkan data yang lebih bersih dan terstruktur, sehingga hanya kata-kata yang relevan dan penting saja yang tersisa untuk dianalisis lebih lanjut.

Tabel 3.2 Contoh Proses *Data Cleaning*

Sebelum <i>Data Cleaning</i>	DAMPAK EFISIENSI ANGGARAN INI ! @prabowo Dampaknya ke Rakyat yang butuh cari biaya hidup. Ttp yg mengherankan justru terus memperbanyak mengangkat STAFSUS yang TIDAK ADA URGENSINYA https://t.co/Y2OjtNKH9N
Sesudah <i>Data Cleaning</i>	DAMPAK EFISIENSI ANGGARAN INI Dampaknya ke Rakyat yang butuh cari biaya hidup Ttp yg mengherankan justru terus memperbanyak mengangkat STAFSUS yang TIDAK ADA URGENSINYA

3.2.3.1 *Case Folding*

Case folding adalah proses mengubah semua huruf dalam teks menjadi huruf kecil (*lowercase*) agar analisis lebih konsisten dan tidak terpengaruh oleh perbedaan kapitalisasi huruf. Seperti contohnya : “Efisiensi Anggaran” menjadi “efisiensi anggaran”.

Tabel 3.3 Contoh Proses *Case Folding*

Sebelum <i>Case Folding</i>	DAMPAK EFISIENSI ANGGARAN INI Dampaknya ke Rakyat yang butuh cari biaya hidup Ttp yg mengherankan justru terus memperbanyak mengangkat STAFSUS yang TIDAK ADA URGENSINYA
Sesudah <i>Case Folding</i>	dampak efisiensi anggaran ini dampaknya ke rakyat yang butuh cari biaya hidup ttp yg mengherankan justru terus memperbanyak mengangkat stafsus yang tidak ada urgensinya

3.2.3.2 *Normalisasi*

Normalisasi dalam *preprocessing text* merupakan perubahan kata *slank*. *Slank* merupakan kata dalam bentuk bahasa yang dipakai oleh kelompok tertentu untuk berkomunikasi secara santai dan tidak formal. *Slang* dapat terbentuk melalui singkatan, penggunaan kata dengan makna berbeda dari arti aslinya, atau penggabungan kata-kata yang tidak biasa. Dalam bidang pengolahan bahasa alami

(*Natural Language Processing*), *slang* sering kali memiliki arti yang kurang jelas atau ambigu, serta maknanya bisa berbeda tergantung pada konteks dan budaya subkelompok tertentu. Penggunaan kata slang ini banyak ditemukan pada media sosial. Misalkan kata “ttp” yang merupakan singkatan dari “tetapi”.

Tabel 3.4 Contoh Proses Normalisasi

Sebelum Normalisasi	dampak efisiensi anggaran ini dampaknya ke rakyat yang butuh cari biaya hidup ttp yg mengherankan justru terus memperbanyak mengangkat stafsus yang tidak ada urgensinya
Sesudah Normalisasi	dampak efisiensi anggaran ini dampaknya ke rakyat yang butuh cari biaya hidup tetapi yang mengherankan justru terus memperbanyak mengangkat stafsus yang tidak ada urgensinya

3.2.3.3 *Stopwords Removal*

Stopwords removal adalah langkah untuk menghapus kata-kata yang sering muncul dalam teks, namun tidak memberikan informasi penting dalam analisis. Kata-kata seperti "dan", "atau", "dengan", "yang", dan sejenisnya biasanya dihapus.

Tabel 3.5 Contoh Proses *Stopwords Removal*

Sebelum <i>Stopwords Removal</i>	dampak efisiensi anggaran ini dampaknya ke rakyat yang butuh cari biaya hidup ttp yg mengherankan justru terus memperbanyak mengangkat stafsus yang tidak ada urgensinya
Sesudah <i>Stopwords Removal</i>	dampak efisiensi anggaran dampaknya rakyat butuh cari biaya hidup mengherankan justru terus memperbanyak mengangkat stafsus tidak urgensi

3.2.3.4 *Stemming*

Stemming adalah proses mengubah kata-kata menjadi bentuk dasarnya (akar kata), sehingga kata-kata turunan seperti "ubah", "merubah", dan "diubah" akan menjadi "ubah".

Tabel 3.6 Contoh Proses *Stemming*

Sebelum <i>Stemming</i>	dampak efisiensi anggaran dampaknya rakyat butuh cari biaya hidup mengherankan justru terus memperbanyak mengangkat stafsus tidak urgensi
Sesudah <i>Stemming</i>	dampak efisiensi anggaran dampak rakyat butuh cari biaya hidup heran justru terus banyak angkat stafsus tidak urgensi

3.2.3.5 *Tokenizing*

Tokenizing adalah proses memecah teks menjadi potongan-potongan kecil, yang disebut "token", biasanya berupa kata atau kalimat. Proses ini memudahkan pengolahan teks lebih lanjut dalam analisis. *Tokenizing* sangat penting dalam tahap *preprocessing* data pada penelitian ini karena peneliti menggunakan pembobotan TF-IDF, yang akan lebih mudah diterapkan setelah data diubah menjadi beberapa token.

Tabel 3.7 Contoh Proses *Tokenizing*

Sebelum <i>Tokenizing</i>	dampak efisiensi anggaran dampak rakyat butuh cari biaya hidup heran justru terus banyak angkat stafsus tidak urgensi
Sesudah <i>Tokenizing</i>	['dampak', 'efisiensi', 'anggaran', 'dampak', 'rakyat', 'butuh', 'cari', 'biaya', 'hidup', 'heran', 'justru', 'terus', 'banyak', 'angkat', 'stafsus', 'tidak', 'urgensi']

3.2.4 Perhitungan TF-IDF

Tahap selanjutnya yaitu perhitungan pembobotan kata dengan metode TF-IDF. Data yang sudah memasuki proses *preprocessing* akan menjadi beberapa token. Sehingga dalam proses perhitungan TF-IDF akan lebih mudah dilakukan. Setiap kata yang ada pada data akan diberikan bobot untuk bisa diolah ke dalam sebuah model. Hal tersebut dilakukan agar model dapat mengetahui kata yang paling relevan dan paling informatif dalam sebuah sentimen.

Dalam penelitian ini setiap kata yang muncul dalam tweet dianggap sebagai entitas yang akan digunakan dalam analisis. Dengan kata lain, vocabulary yang dibangun berfungsi sebagai "kamus" yang mencakup seluruh kata yang digunakan dalam corpus data, dan setiap kata diberikan indeks unik. Sehingga akan ada dua kali perhitungan untuk menentukan pembobotan setiap kata, perhitungan TF dan Perhitungan IDF. Perhitungan TF (*Term Frequency*) sendiri merupakan

perhitungan yang dilakukan untuk seberapa sering sebuah kata muncul dalam sebuah dokumen tertentu. Rumus untuk perhitungan TF ada pada persamaan 2.1

Setelah didapatkan bobot dari TF Selanjutnya, untuk memberikan bobot yang lebih besar pada kata-kata yang jarang muncul di banyak dokumen, kita menghitung *Inverse Document Frequency* (IDF). Perhitungan ini untuk mengukur seberapa penting sebuah kata dalam seluruh corpus (kumpulan dokumen). Semakin jarang kata tersebut muncul dalam dokumen-dokumen lain, semakin besar bobotnya. Rumus untuk menghitung IDF ada pada persamaan 2.2.

Bobot dari hasil perhitungan TF dan IDF untuk setiap kata telah dilakukan. Langkah selanjutnya adalah mengalikan keduanya untuk menghasilkan TF-IDF dengan menggunakan rumus persamaan 2.3. Proses ini memberikan bobot lebih pada kata-kata yang sering muncul dalam satu dokumen namun jarang muncul di dokumen lain, sehingga kata-kata yang lebih spesifik dan informatif mendapatkan nilai lebih tinggi. Dengan mengalikan TF dan IDF untuk setiap kata dalam setiap dokumen, terbentuklah matriks TF-IDF.

a. *Reshape Data Matriks*

Hasil perhitungan bobot menggunakan metode TF-IDF berupa matriks 2 dimensi yang berisi bobot untuk setiap kata dalam dokumen atau tweet. Matriks TF-IDF ini pada dasarnya menggambarkan representasi numerik dari setiap kata dalam corpus yang telah melalui proses tokenisasi, dengan setiap elemen dalam matriks tersebut merepresentasikan pentingnya kata pada suatu dokumen.

Namun, untuk dapat digunakan dalam model *Long Short-Term Memory* (LSTM), data hasil perhitungan TF-IDF harus diubah terlebih dahulu menjadi format yang sesuai dengan kebutuhan input model. Model LSTM membutuhkan input dalam format 3 dimensi dengan urutan sebagai berikut: [samples, time steps, features], di mana:

- *Samples* adalah jumlah data tweet yang digunakan.
- *Time Steps* adalah jumlah langkah waktu (*timestep*) dalam urutan input.
- *features* adalah jumlah fitur yang ada di setiap langkah waktu atau hasil fitur TF-IDF dari setiap kata data *tweet*.

Pada umumnya, hasil TF-IDF menghasilkan matriks 2 dimensi dengan bentuk [*samples, features*], di mana setiap baris mewakili satu data tweet dan setiap kolom mewakili bobot TF-IDF untuk kata-kata yang ada dalam dokumen. Agar dapat diterima oleh model LSTM, matriks 2 dimensi ini perlu diubah menjadi matriks 3 dimensi dengan menambahkan dimensi untuk *time steps*. Dalam penelitian ini, digunakan 1 *timestep* sebagai langkah waktu. Matriks ini menggambarkan representasi numerik yang lebih tepat tentang pentingnya kata-kata dalam konteks seluruh corpus. Pentingnya perhitungan pembobotan kata agar bisa digunakan untuk pembuatan model klasifikasi dengan tujuan memahami pola atau hubungan dalam data teks secara lebih mendalam. Misalnya, matriks [3, 1, 4] maka bentuk dari *matriks* yang digunakan sebagai inputan yaitu:

$$\begin{bmatrix} [0.2 & 0.5 & 0.3 & 0.1] \\ [0.4 & 0.1 & 0.6 & 0.2] \\ [0.3 & 0.3 & 0.2 & 0.5] \end{bmatrix}$$

Dari matriks diatas samples merupakan 3 baris yang mana menunjukkan 3 data *tweet* yang digunakan. Lalu *features* dari data tersebut berjumlah 4 dengan setiap kata memiliki bobot masing” pada setiap data. Sehingga contoh matriks seperti ini yang akan digunakan sebagai inputan model.

b. *One-Hot Encoding*

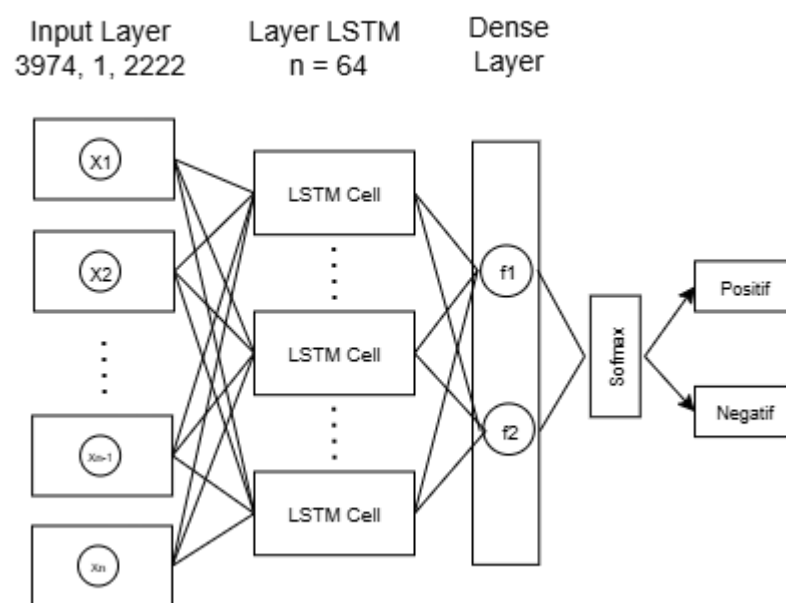
Tidak hanya data *tweet* saja yang dijadikan numerik tetapi label kategori juga diubah format yang dapat diproses oleh algoritma pembelajaran mesin. Pada tahap ini, Label Encoding digunakan untuk mengonversi label kelas berbentuk kategori (misalnya "positif" dan "negatif") menjadi angka numerik. Sebagai contoh, label "positif" bisa diubah menjadi 0 dan "negatif" menjadi 1. Penelitian ini menggunakan *One-Hot Encoding*. Setiap label kategori diubah menjadi vektor biner, di mana hanya satu elemen dalam vektor yang bernilai 1, sementara elemen lainnya bernilai 0. Misalnya, jika ada dua kelas, "positif" dan "negatif", maka label "positif" akan diubah menjadi [1, 0], dan label "negatif" menjadi [0, 1].

3.2.5 Implementasi *Long Short-Term Memory* (LSTM)

Data teks yang sudah diubah menjadi data numerik selanjutnya akan diolah ke dalam model yang digunakan pada penelitian ini, yaitu *Long Short-Term Memory* (LSTM). LSTM merupakan jenis jaringan saraf tiruan yang dirancang untuk menangani data urutan atau *sequence*, sehingga sangat efektif dalam memproses dan memprediksi data berbasis waktu atau urutan kata dalam teks. LSTM memiliki kemampuan untuk mengingat informasi dalam jangka panjang, yang membedakannya dari jaringan saraf konvensional. Dalam arsitektur LSTM,

terdapat beberapa komponen penting yang bekerja secara bersama-sama untuk memproses informasi dalam urutan data.

Pada penelitian yang sedang dilakukan, arsitektur LSTM yang digunakan mengacu pada desain penelitian yang telah dilakukan oleh (Musfiroh et al., 2024) berikut akan ditunjukkan arsitektur tersebut pada gambar 3.4.



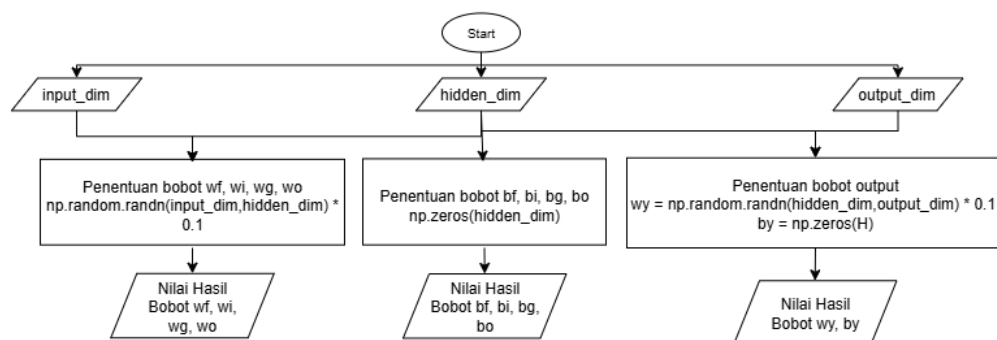
Gambar 3. 4 Arsitektur LSTM

Pada penelitian yang dilakukan sesuai dengan gambar 3.4 arsitektur model dimulai dengan input layer yang menerima teks yang telah diproses menjadi representasi numerik menggunakan TF-IDF, menghasilkan vektor dengan panjang yang sesuai dengan jumlah kata dalam *vocabulary* dengan hasil dimensi TF-IDF yaitu 3974, 1, 2222. Vektor tersebut kemudian diteruskan ke LSTM layer yang terdiri dari 64 unit, yang berfungsi untuk memproses urutan kata dalam teks dan menangkap informasi kontekstual pada setiap timestep. Selanjutnya, output dari

LSTM layer, yaitu hidden state terakhir, diteruskan ke *dense layer* dengan dua unit, yang mengubahnya menjadi dua nilai output yang mewakili probabilitas untuk kelas Positif dan Negatif. Terakhir, *softmax activation* diterapkan pada output dari *dense layer* untuk mengonversi nilai tersebut menjadi *probabilitas*, memastikan bahwa hasilnya berada dalam rentang 0 hingga 1, dan jumlah total *probabilitas* adalah 1. Model ini digunakan untuk klasifikasi sentimen biner berdasarkan teks yang diberikan.

a. Inisialisasi Nilai Bobot dan Bias

Proses pengolahan data menggunakan metode LSTM pada penelitian ini dimulai dari menginisialisasikan beberapa nilai yang nantinya akan di gunakan pada proses *forward* dan *backpropagation* dalam implementasi LSTM. Berikut akan ditampilkan flowchart terkait inisialisasi yang dilakukan pada gambar 3.5.



Gambar 3. 5 Flowchart Inisialisasi Nilai

Pada gambar 3.5 dilakukan adalah inisialisasi parameter-parameter yang diperlukan untuk membangun model. Pada bagian ini, beberapa matriks bobot dan bias diinisialisasi untuk mendukung operasi LSTM. Dimensi input, hidden, dan

output ditentukan oleh parameter, sedangkan nilai *learning rate* dan *dropout rate* juga diatur.

Model LSTM terdiri dari empat gerbang utama, yaitu *forget gate* (f), *input gate* (i), *candidate cell state* (g), dan *output gate* (o). Masing-masing gerbang ini memiliki bobot yang menghubungkan input ke *hidden state*. Matriks bobot yang diinisialisasi untuk masing-masing gerbang adalah W_f , W_i , W_g , W_o untuk hubungan *input-to-hidden*. Bobot-bobot ini diinisialisasi secara acak dengan faktor pembobotan 0.1 agar tidak memiliki nilai awal yang terlalu besar, yang dapat mengganggu proses pelatihan.

Selain itu, bias untuk setiap gerbang juga diinisialisasi dengan nilai nol. Bias-bias ini termasuk b_f untuk *forget gate*, b_i untuk *input gate*, b_g untuk *candidate cell state*, dan b_o untuk *output gate*. Bias ini diperlukan untuk menyesuaikan hasil output dari setiap gerbang, yang membantu model untuk lebih fleksibel dalam mempelajari pola dari data.

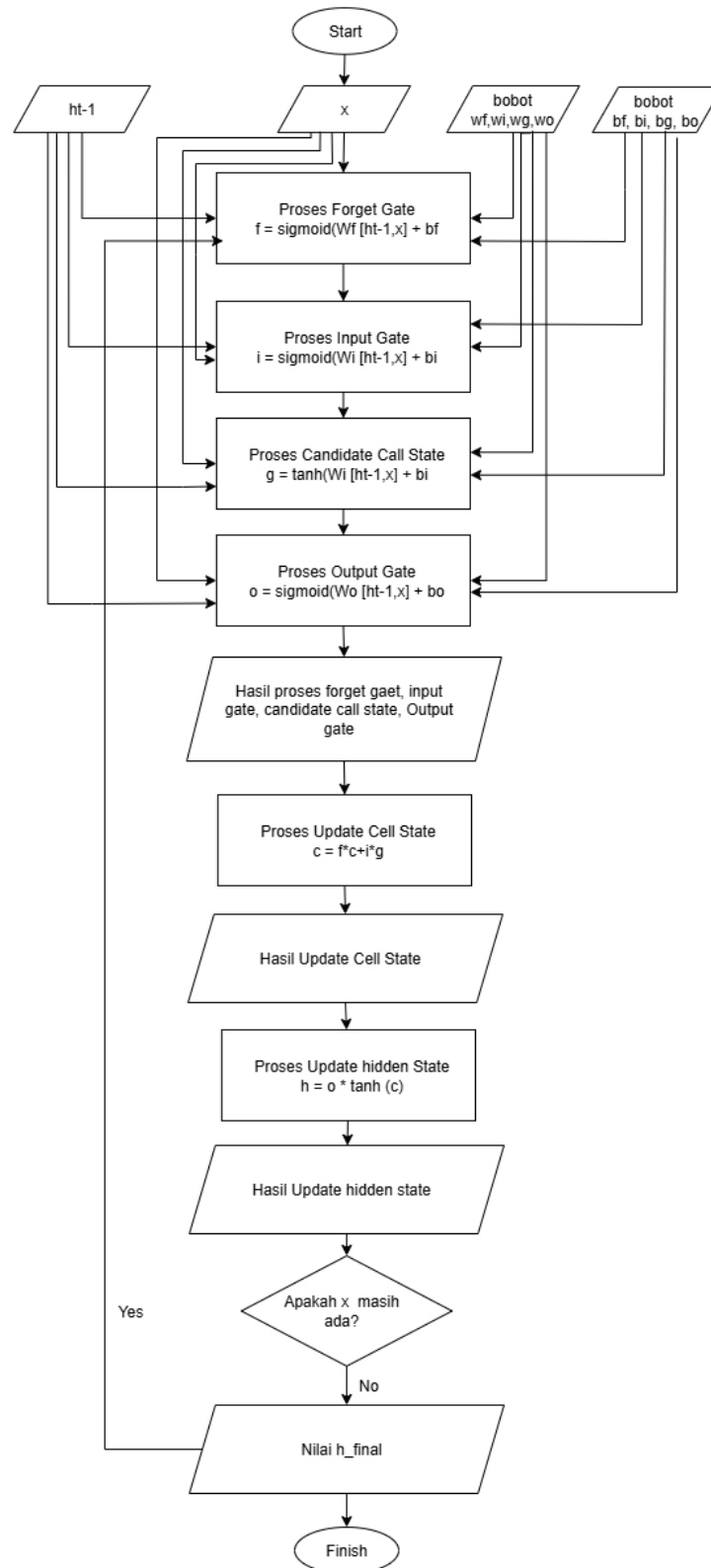
Selain bobot dan bias untuk gerbang-gerbang LSTM, model ini juga menginisialisasi bobot untuk output layer (W_y) dan bias (b_y). Bobot W_y digunakan untuk mentransformasikan hasil dari *hidden state* terakhir menjadi output model, sedangkan bias b_y diterapkan pada output tersebut.

Seluruh parameter ini digunakan dalam proses *forward*, di mana input diproses melalui LSTM untuk menghasilkan output. Proses ini dimulai dengan perhitungan gerbang-gerbang LSTM untuk setiap langkah waktu, diikuti dengan pembaruan status sel (*cell state*) dan *hidden state* berdasarkan nilai gerbang tersebut.

b. *Forward*

Forward dalam LSTM adalah proses pengolahan data urutan dengan menggunakan tiga gerbang utama: forget gate, input gate, dan output gate. Forget gate memutuskan informasi mana yang perlu dilupakan, input gate mengontrol informasi baru yang akan ditambahkan, dan output gate menghasilkan output berdasarkan *cell state* dan *hidden state* yang diperbarui. Ketiga gerbang ini bekerja sama untuk memastikan model dapat mempelajari ketergantungan jangka panjang dalam data urutan, memungkinkan jaringan untuk memproses data panjang tanpa kehilangan informasi penting.

Proses ini berlanjut untuk setiap elemen dalam urutan input, dengan *hidden state* dan *cell state* yang diperbarui di setiap langkah. Model mampu menyimpan informasi relevan dan melupakan yang tidak perlu, sehingga meningkatkan akurasi prediksi. Selain itu, LSTM mengatasi masalah *vanishing gradient* yang sering ditemui dalam jaringan saraf tradisional, serta efektif dalam menangani ketergantungan jangka panjang. Pendekatan ini sangat berguna dalam aplikasi seperti pemrosesan bahasa alami, analisis sentimen, dan prediksi deret waktu. Berikut *flowchart forward* pada penelitian ini yang menggambarkan alur detail dari setiap langkah dalam proses implementasi LSTM.



Gambar 3. 6 Flowchart Implementasi LSTM

Diterapkan fungsi *forward pass*, input yang digunakan pada proses ini adalah h_{t-1} , x , bobot (w) secara keseluruhan yang sudah diinisialisasikan serta bobot bias (b). Kemudian, untuk setiap elemen dalam urutan input, nilai untuk empat gerbang utama LSTM dihitung. Keempat gerbang tersebut adalah *forget gate* (f), *input gate* (i), *candidate cell state* (g), dan *output gate* (o). Setiap gerbang dihitung menggunakan operasi matriks, dengan fungsi aktivasi yang berbeda untuk masing-masing gerbang. Untuk fungsi sigmoid dihitung pada tiga gate yaitu *forget gate* (f), *input gate* (i), dan *output gate* (o). Sedangkan fungsi tanh digunakan pada gate *candidate cell state* (g). Rumus untuk perhitungan setiap gerbang adalah sebagai berikut:

$$f = \sigma(W_f[h_{t-1}, x] + b_f) \quad (3.1)$$

$$i = \sigma(W_i[h_{t-1}, x] + b_i) \quad (3.2)$$

$$g = \tanh(W_g[h_{t-1}, x] + b_g) \quad (3.3)$$

$$o = \sigma(W_o[h_{t-1}, x] + b_o) \quad (3.4)$$

Setelah menghitung nilai untuk setiap gerbang, *cell state* (c) dan *hidden state* (h) diperbarui. Pembaruan *cell state* (c) dihitung dengan menggabungkan hasil dari *forget gate* dan *input gate* untuk memodifikasi status sel sebelumnya dan menambahkan informasi baru dari *candidate cell state*. *Hidden state* (h) dihitung dengan menggunakan *output gate* yang mengontrol informasi yang diteruskan dari status sel ke *hidden state*. Rumus untuk pembaruan cell state dan hidden state adalah sebagai berikut:

$$c = f \cdot c + i + g \quad (3.5)$$

$$h = o. \tanh(c) \quad (3.6)$$

Hasil perhitungan status tersembunyi dan status sel ini disimpan di dalam x hingga selesai diproses sehingga hasil akhirnya diinisialisasikan h_{final} . Setelah *forward* dilakukan untuk menghasilkan prediksi output, proses *backpropagation* dimulai untuk menghitung *gradien error* dari fungsi *loss* terhadap semua parameter dalam model. Dalam LSTM, proses ini mengakumulasi gradien secara berurutan, dimulai dari output layer hingga input layer, yang dikenal dengan nama *backpropagation through time* (BPTT). Gradien yang dihitung ini mencakup semua parameter, seperti bobot dan bias di setiap gerbang LSTM (*forget gate*, *input gate*, *output gate*, dan *candidate cell state*). Gradien tersebut kemudian digunakan untuk memperbarui bobot dan bias menggunakan algoritma *gradient descent*, yang memungkinkan model untuk meminimalkan kesalahan prediksi secara iteratif. Pembaruan bobot dan bias dilakukan dengan mengurangi nilai parameter berdasarkan gradien yang dihitung, yang dikalikan dengan *learning rate*, untuk memperbaiki kinerja model pada iterasi berikutnya.

BAB IV

HASIL DAN PEMBAHASAN

4.1 Skenario Pengujian

Penelitian ini menguji sistem melalui serangkaian langkah yang dirancang untuk memperoleh kinerja model secara efisien. Dengan kategori klasifikasi data, evaluasi akan dilakukan untuk menilai kemampuan model dapat mengklasifikasikan topik penelitian yaitu tanggapan masyarakat terkait efisiensi anggaran 2025. Setiap langkah pengujian melibatkan variasi *hyperparameter* model yang akan digunakan dalam proses pelatihan yang berguna untuk mendapatkan hasil yang optimal. Skenario uji coba yang diterapkan dalam penelitian ini adalah sebagai berikut.

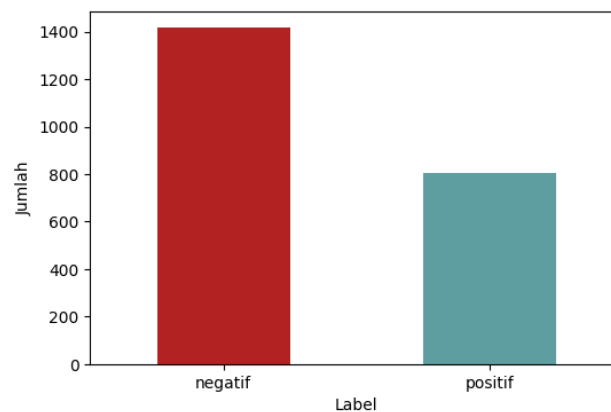
4.1.1 Data Uji

Dataset yang digunakan diperoleh dari platform *Kaggle* dan terdiri atas 2.222 data teks yang diberikan label ke dalam dua kategori, yaitu kelas positif dan kelas negatif. Sebelum digunakan dalam pemodelan, data terlebih dahulu melalui tahapan *preprocessing*, seperti pembersihan teks, serta pemberian bobot menggunakan metode *Term Frequency–Inverse Document Frequency* (TF-IDF). Selanjutnya, data dibagi menjadi data latih dan data uji dengan tiga variasi rasio pembagian, yakni 70:30, 80:20, dan 90:10. Informasi terkait jumlah data uji pada setiap kelas akan ditampilkan pada Tabel 4.1.

Tabel 4.1 Jumlah Data Uji

Label	Jumlah	Persentase
Positif	806	36,27%
Negatif	1416	63,73%

Tabel 4.1 menunjukkan bahwa jumlah tanggapan negatif dalam dataset lebih dominan dibandingkan tanggapan positif. Sebagian besar masyarakat menyatakan ketidaksetujuan terhadap kebijakan efisiensi anggaran 2025 karena dinilai menyebabkan pengurangan fasilitas publik yang merugikan. Namun, terdapat pula masyarakat yang mendukung kebijakan tersebut karena memahami pentingnya efisiensi anggaran dalam menjaga keberlanjutan pembangunan. Gambar 4.1 merupakan grafik perbandingan data kategori jumlah label.



Gambar 4.1 Grafik Perbandingan Data Kategori

Pengujian dalam penelitian ini dilakukan dengan membagi dataset ke dalam tiga variasi rasio dengan tujuan untuk mengetahui rasio pembagian data yang memberikan performa model terbaik. Pada rasio 70:30, dari total 2.222 data, sebanyak 1.555 data digunakan sebagai data latih dan 667 sebagai data uji. Pada rasio 80:20, terdapat 1.778 data latih dan 444 data uji, sedangkan pada rasio 90:10, sebanyak 1.999 data digunakan untuk pelatihan dan 223 untuk pengujian. Variasi rasio ini bertujuan untuk mengamati pengaruh proporsi data latih terhadap akurasi model dalam skenario klasifikasi data yang tidak seimbang. Pada setiap pembagian, distribusi kelas positif dan negatif dijaga agar tetap seimbang. Namun, karena

keterbatasan jumlah data, jumlah data pada kelas positif tetap lebih sedikit dibandingkan kelas negatif. Selain itu, *random seed* juga ditentukan untuk memastikan proses pengacakan data dilakukan secara konsisten. Rincian pembagian data pada masing-masing rasio ditampilkan pada Tabel 4.2.

Tabel 4.2 Pembagian Rasio Data

<i>Rasio</i>	<i>Data Training</i>	<i>Data Testing</i>	<i>Random Seed</i>
70 : 30	1555	667	42
80 : 20	1778	444	42
90 : 10	1999	223	42

4.1.2 Mengukur performa sistem

Pengujian sistem yang telah dilakukan dalam penelitian ini dilanjutkan dengan tahap pengukuran performa model. Data yang telah melalui proses pembagian rasio serta percobaan variasi *hyperparameter* kemudian dievaluasi untuk mengetahui tingkat akurasi model. Evaluasi dilakukan menggunakan *confusion matrix* untuk membandingkan hasil prediksi model LSTM dengan data aktual yang telah dilabeli. Berikut akan ditampilkan pada gambar 4.2 evaluasi *confusion matrix*.

True Labels	0 (Negatif)	TN	FP
	1 (Positif)	FN	TP
		0 (Negatif)	1 (Positif)
		Actual Values	

Gambar 4.2 Visualisasi *Confusion Matrix*

Pada gambar 4.2 tersebut juga terdapat empat istilah yang merupakan representasi hasil proses klasifikasi, istilah tersebut akan dijelaskan sebagai berikut:

1. *True Positive* (TP): Jumlah tweet yang benar-benar memiliki sentimen positif dan diprediksi positif oleh model
2. *True Negative* (TN): Jumlah tweet yang benar-benar memiliki sentimen negatif dan diprediksi negatif oleh model.
3. *False Positive* (FP): Jumlah tweet yang sebenarnya memiliki sentimen negatif tetapi diprediksi positif oleh model (juga disebut sebagai "*Type I Error*").
4. *False Negative* (FN): Jumlah tweet yang sebenarnya memiliki sentimen positif tetapi diprediksi negatif oleh model (juga disebut sebagai "*Type II Error*").

Empat istilah yang ada pada *Confusion matrix* digunakan untuk menghitung nilai akurasi, presicion, *recall* dan F1-score. Akurasi merupakan metrik yang digunakan untuk mengukur sejauh mana model berhasil memprediksi dengan benar. Perhitungan akurasi dapat dihitung menggunakan rumus yang tercantum pada Persamaan 4.1.

$$Akurasi = \frac{TP + TN}{TP + TN + FP + FN} \times 100\% \quad (4.1)$$

Presisi adalah rasio antara prediksi positif yang benar (*True Positive*) dan seluruh prediksi positif yang dibuat model (*True Positive* + *False Positive*), yang menggambarkan sejauh mana model akurat dalam memprediksi kelas positif.

Perhitungan presisi dapat dihitung menggunakan rumus yang ada pada persamaan 4.2

$$Precision = \frac{TP}{TP + FP} \times 100\% \quad (4.2)$$

Recall adalah rasio antara jumlah prediksi positif yang benar (*True Positive*) dan seluruh data yang sebenarnya positif (*True Positive + False Negative*), yang menunjukkan seberapa baik model dalam menemukan atau mengenali semua data positif. Berikut rumus *recall* pada persamaan 4.3

$$Recall = \frac{TP}{TP + FN} \times 100\% \quad (4.3)$$

F1-score merupakan metrik evaluasi yang menggabungkan recall dan presisi dalam satu nilai rata-rata harmonis, memberikan gambaran yang seimbang mengenai kemampuan model untuk mengenali semua data positif (*recall*) dan memprediksi dengan akurat (*presisi*), sangat berguna ketika data yang digunakan tidak seimbang. Persamaan 4.4 merupakan rumus perhitungan F1-score.

$$F1 - Score = \frac{2 \times Precision \times Recall}{Precision + Recall} \times 100\% \quad (4.4)$$

4.2 Pemrosesan Data

Proses ini bertujuan untuk menghapus segala bentuk *noise* atau data kotor pada teks dan menyiapkannya untuk proses klasifikasi. Proses yang dilakukan mencakup *tokenisasi*, *stemming*, *stopword removal*, normalisasi, *case folding*, dan *data cleaning*. Tahapan ini sangat penting untuk memastikan bahwa data yang digunakan dalam model LSTM sudah bersih dan siap untuk dianalisis lebih lanjut.

4.2.1. Data Cleaning

Fungsi data *cleaning* ini bertujuan untuk mempersiapkan teks untuk analisis dengan menghapus elemen yang tidak relevan, seperti username Twitter, URL, angka, dan tanda baca. Selain itu, karakter khusus seperti "+", "_", "-", dan "/" digantikan dengan spasi, hashtag diproses dengan memisahkan kata, dan huruf yang berulang lebih dari dua kali dikurangi. Terakhir, spasi berlebihan dihapus agar teks lebih rapi dan siap digunakan dalam model klasifikasi. Berikut ditampilkan sampel data cleaning pada tabel 4.1

Tabel 4.3 Hasil *Preprocessing Data Cleaning*

<i>Data Tweet</i>	<i>Data Cleaning</i>
Aku sih gak sepenuhnya yakin bahwa efisiensi dgn pangkas memangkas anggaran K/L ekstrim ini cuman utk biayai program MBG ya. Sebab ada utang jatuh tempo sebesar Rp1.350 T yg harus dibayar pda tahun ini. Ini warisan utang Jokowi ke Prabowo. https://t.co/8gUhM6n6x9	Aku sih gak sepenuhnya yakin bahwa efisiensi dgn pangkas memangkas anggaran K L ekstrim ini cuman utk biayai program MBG ya Sebab ada utang jatuh tempo sebesar Rp T yg harus dibayar pda tahun ini Ini warisan utang Jokowi ke Prabowo
@prabu_yudianto Efisiensi anggaran hanya membuat dana is berkurang sedikit masih ada setidaknya 1t untuk di hancurkan ke pos2 budaya infrastruktur monarki. Pemanfaatan dari tahun ke tahun tidak bisa d telusuri laporan auditnya. Jingan	Efisiensi anggaran hanya membuat dana is berkurang sedikit masih ada setidaknya t untuk di hancurkan ke pos budaya infrastruktur monarki Pemanfaatan dari tahun ke tahun tidak bisa d telusuri laporan auditnya Jingan
@susipudjiastuti Efisiensi & program yang memprioritaskan kebutuhan rakyat (bukan pejabat) sesuai prioritas kemendesakannya serta tepat sasaran. Selama ini (paling tidak 10 th terakhir) terkesan penghamburan anggaran markup gila²an dibiarkan agar pejabat senang dan beri dukungan.	Efisiensi amp program yang memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya serta tepat sasaran Selama ini paling tidak th terakhir terkesan penghamburan anggaran markup gila²an dibiarkan agar pejabat senang dan beri dukungan

Terlihat pada tabel, kolom data *tweet* terdapat satu sampel data yaitu “Ini warisan utang Jokowi ke Prabowo. <https://t.co/8gUhM6n6x9>” serta terdapat contoh sampel lagi dengan kalimat “sedikit masih ada setidaknya 1t “. Pada contoh kedua sampel kalimat tersebut terdapat *hyperlink* dan juga angka yang mana dlam *texts*

preprocessing tidak digunakan. Hal tersebut dikarenakan dalam pengolahan *text* hanya kalimat penting dan bermakna yang akan digunakan untuk pengolahan data yang lebih baik. Sehingga, dalam data cleaning digunakan sebagai proses dengan menghapus elemen yang *noise* seperti spasi berlebih atau karakter lainnya. Hasil data nya akan digunakan untuk proses selanjutnya yaitu proses *case folding*.

4.2.2. Case Folding

Proses *case folding* guna mengubah semua teks yang telah dibersihkan untuk dijadikan huruf kecil. Hal ini dilakukan untuk memastikan konsistensi dalam teks, menghindari perbedaan penulisan yang tidak relevan antara huruf besar dan kecil. Setelah penerapan *case folding* pada kolom '*cleaned_text*', hasilnya disimpan dalam kolom baru '*case_folded_text*'. Sampel data ditampilkan pada tabel 4.2

Tabel 4.4 Hasil *Preprocessing Case Folding*

<i>Data Cleaning</i>	<i>Case Folding</i>
Aku sih gak sepenuhnya yakin bahwa efisiensi dgn pangkas memangkas anggaran K L ekstrim ini cuman utk biayai program MBG ya Sebab ada utang jatuh tempo sebesar Rp T yg harus dibayar pda tahun ini Ini warisan utang Jokowi ke Prabowo	aku sih gak sepenuhnya yakin bahwa efisiensi dgn pangkas memangkas anggaran k l ekstrim ini cuman utk biayai program mbg ya sebab ada utang jatuh tempo sebesar rp t yg harus dibayar pda tahun ini ini warisan utang jokowi ke prabowo
Efisiensi anggaran hanya membuat dana is berkurang sedikit masih ada setidaknya t untuk di hamburkan ke pos budaya infrastruktur monarki Pemanfaatan dari tahun ke tahun tidak bisa d telusuri laporan auditnya Jingan	efisiensi anggaran hanya membuat dana is berkurang sedikit masih ada setidaknya t untuk di hamburkan ke pos budaya infrastruktur monarki pemanfaatan dari tahun ke tahun tidak bisa d telusuri laporan auditnya jingan
Efisiensi amp program yang memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya serta tepat sasaran Selama ini paling tidak th terakhir terkesan penghamburan anggaran markup gila²an dibiarkan agar pejabat senang dan beri dukungan	efisiensi amp program yang memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya serta tepat sasaran selama ini paling tidak th terakhir terkesan penghamburan anggaran markup gila²an dibiarkan agar pejabat senang dan beri dukungan

Sampel data yang ditampilkan dalam tabel 4.2 terdapat beberapa kata seperti “mepanRB”, “MBG”, “CPNS”, dll, diganti ke dalam huruf kecil meskipun

singkatan. Proses klasifikasi harus memiliki token yang seragam dimana ketika kata yang sama maka dapat dihitung menjadi satu token. Sehingga misalkan “Anggaran” dengan “anggaran” menjadi satu kata yang sama untuk dihitung. Maka dilakukan proses *case folding* tersebut.

4.2.3. Normalisasi

Data yang sudah di jadikan huruf kecil pada semua data, selanjutnya akan dilakukan pemeriksaan kata per kata dan apabila ada singkatan yang sesuai dalam kamus, singkatan tersebut akan diganti dengan kata lengkapnya. Kamus singkatan tersebut dimuat dari file JSON yang sudah peneliti buat sehingga nantinya program dapat memanggil dan mengetahui kata apa saja yang harus diganti dengan kata aslinya. Setelah semua data dibaca dan dianalisis serta menggantinya ke dalam kata asli maka data disimpan dalam kolom normalisasi data. Berikut akan ditampilkan hasil normalisasi data yang sudah dilakukan pada tabel 4.3.

Tabel 4.5 Hasil *Preprocessing* Normalisasi

Case Folding	Normalisasi Data
aku sih gak sepenuhnya yakin bahwa efisiensi dgn pangkas memangkas anggaran k l ekstrim ini cuman utk biayai program mbg ya sebab ada utang jatuh tempo sebesar rp t yg harus dibayar pda tahun ini ini warisan utang jokowi ke prabowo	aku sih gak sepenuhnya yakin bahwa efisiensi dengan pangkas memangkas anggaran kementerian negara lembaga ekstrim ini cuman untuk biayai program makan bergizi gratis ya sebab ada utang jatuh tempo sebesar rp t yang harus dibayar pda tahun ini ini warisan utang jokowi ke prabowo
efisiensi anggaran hanya membuat dana is berkurang sedikit masih ada setidaknya t untuk di hamburkan ke pos budaya infrastruktur monarki pemanfaatan dari tahun ke tahun tidak bisa d telusuri laporan auditnya jingan	efisiensi anggaran hanya membuat dana is berkurang sedikit masih ada setidaknya t untuk di hamburkan ke pos budaya infrastruktur monarki pemanfaatan dari tahun ke tahun tidak bisa d telusuri laporan auditnya jingan
efisiensi amp program yang memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya serta tepat sasaran selama ini paling tidak th terakhir terkesan penghamburan anggaran markup gila²an dibiarkan agar pejabat senang dan beri dukungan	efisiensi amp program yang memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya serta tepat sasaran selama ini paling tidak tahun terakhir terkesan penghamburan anggaran markup gila²an dibiarkan agar pejabat senang dan beri dukungan

Dari tabel 4.3 terdapat contoh singkatan “mbg”, yang mana terdapat dalam kamus file JSON yang sudah dibuat. Maka dalam proses normalisasi data kata tersebut diganti dengan “makan bergizi gratis”. Sehingga nantinya setiap kata pada singkatan memiliki makna atau nilai tersendiri dalam proses perhitungannya.

4.2.4. Stopword Removal

Setiap kata dalam sebuah kalimat tentunya tidak semuanya memiliki makna. Belum tentu kata yang sering muncul akan memberikan kesan yang merujuk ke dalam klasifikasi positif dan negatif. Seperti halnya kata kata yang berada dalam list stopwords, kata “yang”, “di”, “dan”, dan lainnya akan dihilangkan dalam teks. Kamus stopwords dalam penelitian ini menggunakan kamus sastraawi. Peneliti juga menambahkan beberapa list daftar kata yang dihilangkan seperti "amp", "yntkts", dan "dengan", karena dirasa tidak memiliki makna yang mendalam. Setelah proses penghapusan stopwords, teks yang tersisa akan disimpan dalam kolom baru, siap untuk dianalisis lebih lanjut. Tabel 4.6 akan menampilkan hasil contoh sampel datanya.

Tabel 4.6 Hasil *Preprocessing Stopword Removal*

Normalisasi Data	Stopword Removal
aku sih gak sepenuhnya yakin bahwa efisiensi dengan pangkas memangkas anggaran kementerian negara lembaga ekstrim ini cuman untuk biayai program makan bergizi gratis ya sebab ada utang jatuh tempo sebesar rp t yang harus dibayar pda tahun ini ini warisan utang jokowi ke prabowo	aku sih gak sepenuhnya yakin efisiensi pangkas memangkas anggaran kementerian negara lembaga ekstrim cuman biayai program makan bergizi gratis utang jatuh tempo sebesar rp t dibayar pda tahun warisan utang jokowi prabowo
by the way anggaran efisiensi itu gak sama dengan anggaran cpns ya selama belum ada info dari menpanrb pres kurang kurangilah fear mongering kayak gini	by the way anggaran efisiensi gak sama anggaran cpns selama info menpanrb pres kurang kurangilah fear mongering kayak gini
efisiensi amp program yang memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya serta tepat sasaran selama ini paling tidak tahun terakhir terkesan penghamburan anggaran markup gila²an	efisiensi program memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya tepat sasaran selama paling tahun terakhir terkesan penghamburan

Normalisasi Data	Stopword Removal
dibiarkan agar pejabat senang dan beri dukungan	anggaran markup gila ² an dibiarkan pejabat senang beri dukungan

Data yang ada pada sampel tabel 4.4 terdapat sebuah kalimat “*by the way* anggaran efisiensi itu gak sama dengan anggaran cpns ya selama belum ada info dari menpanrb pres kurang kurangnya fear mongering kayak gini “. Dalam kalimat tersebut terdapat beberapa kata yang dihapus seperti “dengan”, “ya”, “belum”, “ada”. Hal itu dikarenakan kata kata tersebut dalam *sastrawi* menjadi list yang harus dihapus. Terdapat juga kata dengan sebagai kata pengecualian yang sudah ditentukan oleh penulis untuk dihilangkan.

4.2.5. Stemming

Dalam sebuah *tweet* terkadang memiliki kata yang terdapat sebuah imbuhan, seperti me-mbangun. Data yang digunakan memiliki banyak variasi kata. Tetapi ketika kata dikembalikan dalam bentuk dasar maka dapat diketahui bahwasannya kata tersebut saling memiliki makna dengan kata lainnya. Proses *stemming* bertujuan untuk mengembalikan kata ke bentuk dasarnya, menghapus imbuhan yang ada pada kata tersebut. Tahap *stemming* telah melewati tahap penghapusan *stopwords* kemudian diubah ke bentuk dasar (*root word*). Sehingga kata yang akan diolah dapat mengurangi variasi kata dalam data yang dapat mempengaruhi hasil analisis.

Tabel 4.7 Hasil *Preprocessing Stemming*

Stopword Removal	Stemming
aku sih gak sepenuhnya yakin efisiensi pangkas memangkas anggaran kementerian negara lembaga ekstrim cuman biayai program makan bergizi gratis utang jatuh tempo sebesar rp t dibayar pda tahun warisan utang jokowi prabowo	aku sih gak sepenuh yakin efisiensi pangkas mangkas anggar tri negara lembaga ekstrim cuman biaya program makan gizi gratis utang jatuh tempo besar rp t bayar pda tahun waris utang jokowi prabowo

Stopword Removal	Stemming
efisiensi anggaran membuat dana is berkurang sedikit t hamburkan pos budaya infrastruktur monarki pemanfaatan tahun tahun d telusuri laporan auditnya jingan	efisiensi anggar buat dana is kurang sedikit t hambur pos budaya infrastruktur monarki manfaat tahun tahun d telusur lapor audit jingan
efisiensi program memprioritaskan kebutuhan rakyat bukan pejabat sesuai prioritas kemendesakannya tepat sasaran selama paling tahun terakhir terkesan penghamburan anggaran markup gila ² an dibiarkan pejabat senang beri dukungan	efisiensi program prioritas butuh rakyat bukan jabat sesuai prioritas desa tepat sasar lama paling tahun akhir kes hambur anggar markup gila an biar jabat senang beri dukung

Sebagai contoh kata “memprioritaskan” terdapat imbuhan me- dan –an yang mana merupakan bentuk dasar. Kata dasar dari kata tersebut adalah “prioritas”. Sehingga ketika ada kata serupa tetapi berbeda imbuhan maka akan dijadikan satu kata dengan nilai yang sama untuk lebih mempermudah proses.

4.2.6. Tokenizing

Satu kalimat dalam data tweet merupakan satu kesatuan kalimat yang mana harus dilakukan proses tokenisasi untuk memecah teks yang telah melalui tahap *stemming* menjadi unit-unit kata atau token. Hal tersebut juga menjadi langkah yang dapat mempermudah penelitian dikarenakan dalam penelitian ini pembobotan kata yang digunakan yaitu TF-IDF yang mendeteksi setiap token kata memiliki nilai tersendiri. Hasil proses tersebut bisa dilihat dalam tabel 4.8.

Tabel 4.8 Hasil *Preprocessing Tokenizing*

Stemming	Tokenizing
aku sih gak sepenuh yakin efisiensi pangkas mangkas anggar tri negara lembaga ekstrim cuman biaya program makan gizi gratis utang jatuh tempo besar rp t bayar pda tahun waris utang jokowi prabowo	[aku] sih gak sepenuh yakin efisiensi pangkas mangkas anggar tri negara lembaga ekstrim cuman biaya program makan gizi gratis utang jatuh tempo besar rp t bayar pda tahun waris utang jokowi prabowo
efisiensi anggar buat dana is kurang sedikit t hambur pos budaya infrastruktur monarki manfaat tahun tahun d telusur lapor audit jingan	efisiensi anggar buat dana is kurang sedikit t hambur pos budaya infrastruktur monarki manfaat tahun tahun d telusur lapor audit jingan
efisiensi program prioritas butuh rakyat bukan jabat sesuai prioritas desa tepat sasar lama	efisiensi program prioritas butuh rakyat bukan jabat sesuai prioritas desa tepat sasar lama

<i>Stemming</i>	<i>Tokenizing</i>
paling tahun akhir kes hambur anggar markup gila an biar jabat senang beri dukung	paling tahun akhir kes hambur anggar markup gila an biar jabat senang beri dukung

4.3 Hasil Uji Coba dan Evaluasi

Uji coba yang dilakukan dalam penelitian ini adalah menguji kinerja model *Long Short Term Memory* (LSTM) pada berbagai variasi rasio pembagian data pelatihan dan data pengujian. Sebelum masuk ke dalam pengujian data teks di jadikan ke dalam data numerik untuk setiap kata menggunakan metode TF-IDF. Pengujian ini bertujuan untuk menetapkan rasio optimal dalam pembagian data. Rasio yang akan diuji meliputi 70:30, 80:20, dan 90:10. Setiap rasio data akan diuji juga untuk mencari *hyperparameter* seperti *epoch* dan *batch size* yang digunakan. Dengan demikian, penelitian ini bertujuan untuk memperoleh kombinasi parameter terbaik yang dapat meningkatkan kinerja model LSTM dalam klasifikasi tanggapan masyarakat.

4.3.1 Hasil Uji Coba-1

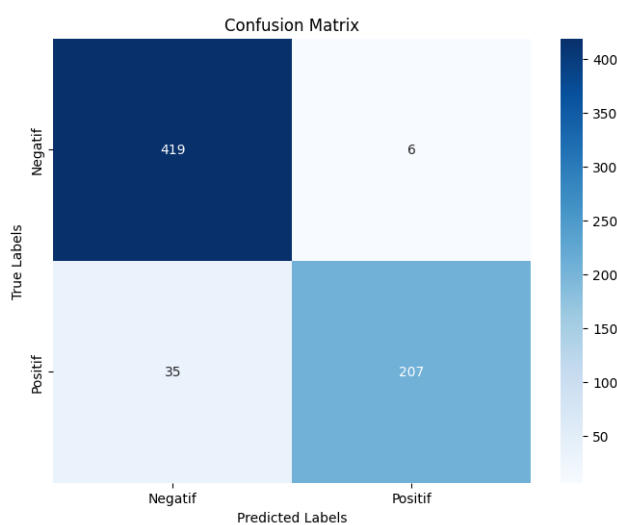
Pada percobaan pertama, dilakukan pengujian dengan variasi rasio data menggunakan rasio 70:30. Hasil pembagian data dalam skenario ini terdiri dari 1555 data untuk pelatihan dan 667 data untuk pengujian. Dalam proses uji coba, diterapkan 425 data negatif dan 242 data positif. Sebelum dilakukan evaluasi, pelatihan model dilakukan dengan menggunakan data pelatihan yang telah diproses dan dibagi sebelumnya. Proses pelatihan ini melibatkan penggunaan *epoch* serta *batch size* untuk mengoptimalkan proses pembelajaran model. Pengujian dilakukan dengan menggunakan tiga *epoch* yaitu 10, 20 dan 30 serta setiap *epoch* akan diuji

dengan nilai *batch size* yaitu 32 dan 64. Hasil pengujian akan ditampilkan pada tabel berikut :

Tabel 4.9 Hasil Pengujian Rasio 70:30

<i>Epoch</i>	<i>Batch Size</i>	<i>Akurasi</i>	<i>Precesion</i>	<i>Recall</i>	<i>F1-Score</i>
10	32	95.20%	97.41%	87.44%	92.16%
20	32	94.90%	98.40%	85.58%	91.54%
30	32	93.70%	92.20%	87.91%	90.00%
10	64	94.60%	94.97%	87.91%	91.30%
20	64	93.55%	91.75%	87.91%	89.79%
30	64	92.95%	90.78%	86.98%	88.84%

Tabel 4.9 menunjukkan hasil pengujian model dengan perbandingan rasio data 70:30 serta kombinasi *epoch* dan *batch size*. Terlihat bahwa peningkatan jumlah epoch dari 10 ke 30 cenderung menurunkan performa model, terutama pada metrik Akurasi, Presisi, dan *F1-Score*. Sebaliknya, penggunaan *batch size* 64 pada *epoch* 10 memberikan performa terbaik, dengan *F1-Score* tertinggi mencapai 90.99%, sementara pada *epoch* 30, *F1-Score* menurun menjadi 88.98%. Gambar 4.3 akan ditunjukkan hasil evaluasi menggunakan *confusion matrix* dari pengujian *epoch* 10 dan *batch size* 64.



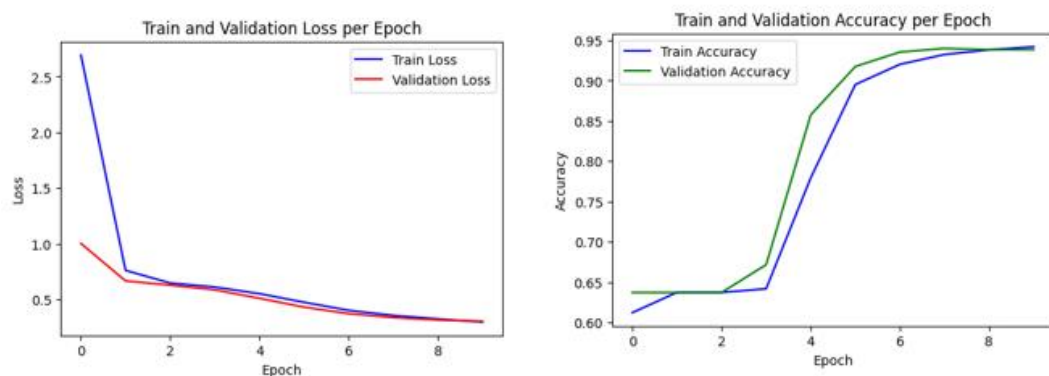
Gambar 4.3 *Confusion Matrix* Epoch 10 dan Batch Size 64 Pada Rasio 70 :30

Pada setiap pelatihan model, epoch adalah jumlah iterasi penuh melalui data pelatihan, sementara batch size menentukan jumlah sampel yang diproses sebelum pembaruan parameter. Hasil *loss* dan akurasi pada data pelatihan dan validasi dicatat per *epoch* untuk memantau kinerja model. Tabel 4.10 berikut akan menampilkan nilai perepoch 10 dengan *batch size* 64

Tabel 4.10 Hasil Pelatihan Model Rasio 70:30

<i>Epoch</i>	<i>Batch Size</i>	<i>Train Loss</i>	<i>Train Accuracy</i>	<i>Validation Loss</i>	<i>Validation Accuracy</i>
1	64	2.6948	0.6122	1.0020	0.6372
2	64	0.7601	0.6373	0.6655	0.6372
3	64	0.6465	0.6373	0.6275	0.6372
4	64	0.6103	0.6418	0.5854	0.6717
5	64	0.5509	0.7794	0.5093	0.8576
6	64	0.4736	0.8952	0.4302	0.9175
7	64	0.4011	0.9203	0.3714	0.9355
8	64	0.3554	0.9325	0.3367	0.9400
9	64	0.3241	0.9383	0.3148	0.9385
10	64	0.2961	0.9421	0.3027	0.9385

Grafik pada gambar 4.4 berikut akan menampilkan gambaran nilai-nilai dari tabel per *epoch*, termasuk loss dan akurasi pada data pelatihan dan validasi, untuk memudahkan analisis kinerja model pada setiap iterasi.



Gambar 4.4 Train and Validation Loss and Accuracy 70:30

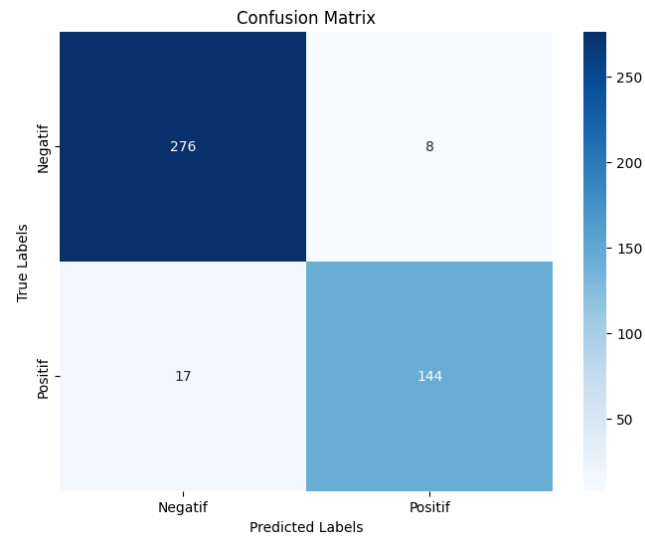
4.3.2 Hasil Uji Coba-2

Pengujian selanjutnya dilakukan dengan menggunakan rasio pembagian data 80:20, di mana data pelatihan terdiri dari 1777 sampel dan data pengujian sebanyak 445 sampel. Pada data pengujian, terdapat 284 data negatif dan 161 data positif. Sebagaimana pada pengujian sebelumnya, pelatihan model dilakukan dengan variasi jumlah epoch, yaitu 10, 20, dan 30, serta dua variasi batch size, yaitu 32 dan 64, untuk setiap iterasi.

Tabel 4.11 Hasil Pengujian Rasio 80:20

<i>Epoch</i>	<i>Batch Size</i>	<i>Akurasi</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
10	32	94.83%	97.41%	84.96%	90.76%
20	32	93.93%	93.44%	85.71%	89.41%
30	32	93.48%	91.27%	86.47%	88.80%
10	64	94.16%	97.35%	82.71%	89.43%
20	64	93.48%	90.62%	87.22%	88.89%
30	64	93.71%	92.00%	86.47%	89.15%

Tabel 4.11 menunjukkan hasil pengujian model dengan rasio pembagian data 80:20 serta kombinasi epoch dan batch size yang berbeda. Hasil pengujian menunjukkan bahwa peningkatan jumlah epoch dari 10 ke 30 menyebabkan penurunan performa model, terutama pada metrik Akurasi, Presisi, dan F1-Score. Model dengan *epoch* 10 dan *batch size* 32 memberikan hasil terbaik, dengan F1-Score tertinggi mencapai 92.01% serta akurasi paling tinggi yaitu 94.38%. Gambar 4.5 akan menampilkan hasil evaluasi menggunakan *confusion matrix* dari pengujian dengan *epoch* 10 dan *batch size* 32.



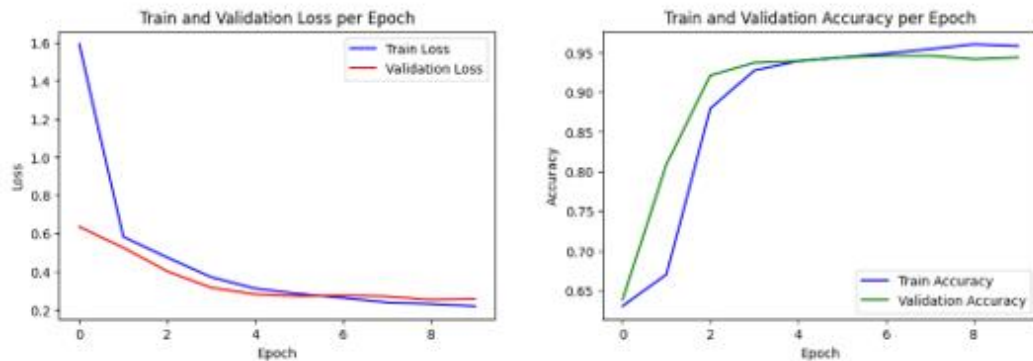
Gambar 4.5 *Confusion Matrix* Epoch 10 dan *Batch Size* 32 Pada Rasio 80:20

Sebelum melakukan evaluasi kinerja model menggunakan *confusion matrix*, model dilatih terlebih dahulu dengan *epoch* dan *batch size* yang telah ditentukan berdasarkan penelitian sebelumnya. Setelah pelatihan selesai, dilakukan pengujian untuk mendapatkan nilai per *epoch*, yang mencakup *train loss*, *train accuracy*, *validation loss*, dan *validation accuracy*. Tabel 4.12 berikut menunjukkan hasil nilai perepoch optimal dari pembagian rasio 80:20.

Tabel 4.12 Hasil Pelatihan Model Rasio 80:20

<i>Epoch</i>	<i>Batch Size</i>	<i>Train Loss</i>	<i>Train Accuracy</i>	<i>Validation Loss</i>	<i>Validation Accuracy</i>
1	32	1.5909	0.6303	0.6349	0.6382
2	32	0.5833	0.6702	0.5258	0.8090
3	32	0.4744	0.8796	0.4025	0.9213
4	32	0.3719	0.9274	0.3166	0.9371
5	32	0.3116	0.9392	0.2812	0.9393
6	32	0.2838	0.9437	0.2719	0.9438
7	32	0.2623	0.9488	0.2754	0.9461
8	32	0.2389	0.9544	0.2702	0.9461
9	32	0.2310	0.9606	0.2523	0.9416
10	32	0.2200	0.9584	0.2580	0.9438

Dengan melihat tabel 4.12 maka dapat digambarkan dengan grafik antara *train loss* dan *train accuracy* serta *validation loss* dan *validation accuracy*. Berikut akan ditampilkan grafik pada gambar 4.6



Gambar 4.6 Train and Validation Loss and Accuracy 80:20

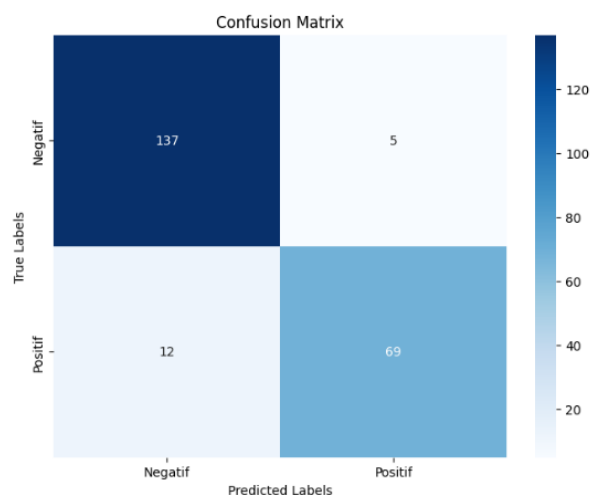
4.3.3 Hasil Uji Coba-3

Pengujian ketiga menggunakan rasio pembagian data 90:10. Seperti pengujian sebelumnya keseluruhan data dibagi menjadi 1999 data pelatihan dan 223 data pengujian. Untuk pengujian digunakan data dengan kategori negatif sebanyak 142 serta kategori positif sebanyak 81. *Hyperparameter* yang digunakan dalam pengujian ketiga ini dengan variasi *epoch* yaitu 10,20, dan 30 serta pada setiap *epoch* diuji dengan *batch size* bernilai 32 dan 64. Hasil pengujian berbagi variasi uji akan ditampilkan pada tabel 4.13.

Tabel 4.13 Hasil Pengujian Rasio 90:10

<i>Epoch</i>	<i>Batch Size</i>	<i>Akurasi</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
10	32	94.62%	93.85%	88.41%	91.04%
20	32	94.17%	92.42%	88.41%	90.37%
30	32	94.62%	95.24%	86.96%	90.20%
10	64	94.17%	95.16%	85.51%	90.51%
20	64	91.93%	94.37%	82.72%	88.16%
30	64	91.48%	90.79%	85.19%	87.90%

Dari Tabel 4.13, diperoleh hasil pengujian dengan pembagian data 90:10. Sama halnya dengan hasil uji coba-2, didapatkan *hyperparameter* optimal pada *epoch* 10 dan *batch size* 32. Akurasi yang diperoleh mencapai nilai tertinggi di antara variasi lainnya, yaitu 92.38%, sementara F1-Score juga menunjukkan keseimbangan yang baik dengan nilai tertinggi 89.03%. Meskipun demikian, setiap peningkatan jumlah epoch tidak menunjukkan peningkatan yang signifikan dalam evaluasi model, melainkan cenderung menurun sedikit demi sedikit. Kedua variasi *batch size* menunjukkan hasil yang serupa, di mana penambahan epoch tidak memberikan peningkatan performa model yang berarti, meskipun nilai-nilai metrik cenderung konstan dan sedikit menurun seiring bertambahnya *epoch*. Evaluasi *confusion matrix* akan dilakukan untuk mengidentifikasi performa model pada Gambar 4.7.



Gambar 4.7 *Confusion Matrix* Epoch 10 dan Batch Size 32 Pada Rasio 90:10

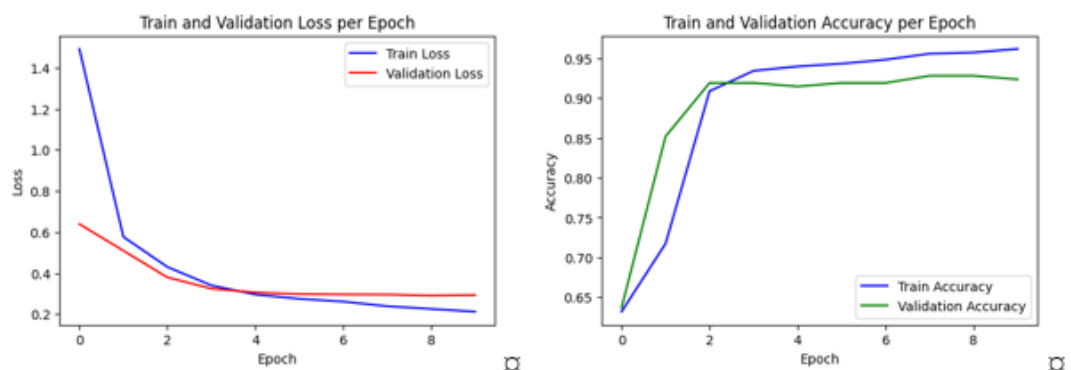
Hasil evaluasi *confusion matrix* yang ditampilkan pada Gambar 4.7 menunjukkan bahwa model berhasil mendeteksi kategori negatif dan positif dengan tingkat akurasi yang tinggi. Model menunjukkan jumlah kesalahan yang relatif

rendah dalam mendeteksi kategori yang salah. Berdasarkan evaluasi ini, dapat dilihat bahwa model mampu mendeteksi kategori dengan baik. Sebelum dilakukan evaluasi model, pelatihan data dilakukan agar model bisa belajar dari data training terlebih dahulu. Saat pelatihan dilakukan peemantau kinerja model. Tabel 4.14 Akan menampilkan hasil pemantauan kinerja model per *epoch* dengan *batchsize* optimal yaitu 32.

Tabel 4.14 Hasil Pelatihan Model Rasio 90:10

<i>Epoch</i>	<i>Batch Size</i>	<i>Train Loss</i>	<i>Train Accuracy</i>	<i>Validation Loss</i>	<i>Validation Accuracy</i>
1	32	1.5909	0.6303	0.6349	0.6382
2	32	0.5833	0.6702	0.5258	0.8090
3	32	0.4744	0.8796	0.4025	0.9213
4	32	0.3719	0.9274	0.3166	0.9371
5	32	0.3116	0.9392	0.2812	0.9393
6	32	0.2838	0.9437	0.2719	0.9438
7	32	0.2623	0.9488	0.2754	0.9461
8	32	0.2389	0.9544	0.2702	0.9461
9	32	0.2310	0.9606	0.2523	0.9416
10	32	0.2200	0.9584	0.2580	0.9438

Dari hasil pelatihan model yang ada pada tabel 4.14 terdapat 10 epoch yang dijalankan. Dari epoch tersebut menghasilkan nilai empat nilai yang tercantum dalam tabel. Nilai tersebut akan di tampilkan pada grafik dengan gambar 4.8.



Gambar 4.8 Train and Validation Loss and Accuracy 90:10

4.4 Pembahasan

Penelitian ini telah melakukan uji coba terhadap 2222 data sentimen masyarakat yang diambil dari dataset *Kaggle* yang berkaitan dengan efisiensi anggaran. Data tersebut harus melalui proses *preprocessing* sebelum dapat diproses oleh sistem. *Preprocessing* ini mencakup beberapa tahapan, seperti data *cleaning*, *case folding*, normalisasi, *stopword removal*, stemming, dan *tokenizing* yang bertujuan agar sistem dapat mengelola data dengan efisien dan menghilangkan noise. Tahap ini sangat penting karena mempengaruhi sejauh mana kinerja sistem dalam mengolah data. Setelah tahap *preprocessing* selesai, langkah pembobotan memiliki peranan penting dalam menentukan hasil klasifikasi sentimen. Pada penelitian ini, pembobotan dilakukan dengan menggunakan metode TF-IDF, yang mengonversi setiap kata atau token menjadi nilai numerik yang bisa diproses oleh sistem.

Dari hasil uji coba pertama, yaitu dengan rasio pembagian data 70:30, diperoleh perbedaan performa pada evaluasi model untuk setiap hyperparameter yang digunakan yaitu *epoch* 10, 20, dan 30 serta *batch size* 32 dan 64 yang telah ditentukan melihat dari penelitian sebelumnya. Pada rasio tersebut, hasil terbaik diperoleh pada saat pengujian *epoch* 10 dengan *batch size* 64. Setelah dilakukan evaluasi pada tabel 4.9 akurasi didapatkan mencapai 93.45% dengan nilai akurasi tertinggi dari variasi pengujian lainnya. Namun, pengujian terbaik tidak hanya didasarkan pada akurasi, melainkan juga didukung oleh nilai *recall* sebesar 85.54%, serta presisi sebesar 97.18% dan F1-Score masing-masing sebesar 90.99%. Dalam

pelatihan model nya juga digunakan *epoch* serta *batch size* dengan *train loss* dan *train accuracy* yang ada pada tabel 4.10

Dilakukan Hasil uji coba kedua yaitu dengan rasio pembagian data 80:20. Skenario uji yang digunakan juga sama dengan uji coba pertama yaitu variasi *hyperparameter*. Berbeda dengan hasil sebelumnya pada uji ini hasil variasi terbaik ada pada *epoch* 10 dengan *batch size* 32. Dari hasil evaluasi yang sudah dilakukan setelah pelatihan model melihat dari tabel 4.11 akurasi tertinggi mencapai 94.38%. tidak hanya itu presisi yang didapatkan juga mencapai nilai tertinggi 94.74%. Disusul dengan *recall* dan *f1-score* sebesar 89.44% dan 92.01%. Dari semua hasil evaluasi tersebut menjadikan variasi tersebut menjadi pengujian terbaik pada pembagian data 80:20.

Uji coba ketiga menggunakan pembagian data dengan rasio 90:10. Pengujian dilakukan dengan berbagai variasi *hyperparameter epoch* serta *batch size*. Hasil dengan uji terbaik ada pada *epoch* 10 dan *batch size* 32 melihat dari evaluasi yang telah dilakukan setelah melatih model. Sama dengan pengujian kedua hasil terbaik berada pada variasi rasio pertama yang ada pada tabel 4.12. Penentuan hasil dengan variasi terbaik dilihat dari keseluruhan nilai evaluasi. Nilai akurasi didapatkan sebesar 92.38% dengan presisi 93.24%. walaupun pada *epoch* 10 *batch size* 64 memiliki presisi yang lebih tinggi tetapi melihat dari nilai evaluasi lainnya seperti *recall* dan *F1-Score* pada *epoch* 10 dan *batch size* 32 menjadi variasi terbaik untuk pengujian pembagian data 90:10.

Setelah melakukan penelitian ketiga uji coba tadi, mengingat bahwa data yang digunakan dalam penelitian ini tidak seimbang, pertimbangan tidak hanya

pada akurasi saja, tetapi juga pada metrik evaluasi lainnya sangat penting. Dengan adanya ketidakseimbangan kelas, model dapat menghasilkan akurasi yang tinggi serta evaluasi yang baik dalam mengklasifikasikan tanggapan masyarakat terkait “Efisiensi Anggaran 2025”. Penambahan pengolahan seperti penggunaan *recall* dan *F1-score* membantu memastikan bahwa model tidak hanya memprediksi kelas mayoritas, tetapi juga mengenali kelas minoritas dengan baik.

Pada pembagian data 70:30 hasil yang didapatkan berbeda dengan rasio 80:20 dan 90:10. Hal tersebut dikarenakan data yang digunakan pada rasio 70:30 dengan jumlah data pelatihan yang sedikit dengan menggunakan batch size 64 menjadi model yang baik memberikan pembaruan parameter yang lebih stabil dan efisien, yang membantu model lebih cepat konvergen dan menghindari *overfitting* pada data yang terbatas. Sedangkan untuk data rasio 80:20, 90:10 menggunakan data pelatihan yang lebih banyak memungkinkan model mempelajari lebih banyak pola tanpa mengalami fluktuasi. Melihat setiap hasil pengujian ketiganya bertambah epoch maka nilai evaluasi yang didapatkan makin menurun menjelaskan bahwa dengan 2222 data serta data yang mungkin tidak terlalu kompleks penggunaan epoch dengan nilai 10 sudah menjadi nilai yang tepat untuk proses pelatihan model. jika nilai *epoch* lebih besar maka nantinya data pelatihan akan makin cepat mengenali data sehingga tidak bisa memprediksi data yang baru.

Hasil penelitian ini menunjukkan pentingnya proses pembersihan *noise* yang menyebabkan adanya kesalahan klasifikasi akibat data *tweet* yang tidak sempurna, seperti adanya singkatan, kata gaul, dan potongan kalimat yang berubah menjadi satu. Proses *cleansing* pada preprocessing menghilangkan link, namun masih ada

kata yang tidak baku atau tidak tepat, seperti kata yang diulang atau disingkat, serta masalah ketidakhadiran spasi dalam kalimat.

4.5 Analisis Konten

Topik yang menjadi objek penelitian ini adalah Kebijakan Efisiensi Anggaran 2025. Data yang digunakan berupa *tweet* yang diambil dari platform aplikasi X, yang banyak membicarakan topik tersebut. Jumlah data yang dianalisis adalah 2222 *tweet* yang terdiri dari 806 *tweet* berlabel positif dan 1416 *tweet* berlabel negatif. Dari hasil pengelompokan *tweet* setiap label, dilakukan analisis untuk melihat kriteria-kriteria yang mendasari setiap label. Peneliti mengidentifikasi masalah dengan memeriksa keseluruhan data *tweet* yang diungkapkan oleh masyarakat. Berdasarkan analisis tersebut, berikut adalah hasil identifikasi masalah yang terkandung dalam setiap kategori label.

4.5.1 Tweet Kategori Positif

Label positif dalam *tweet* yang dianalisis mencerminkan pandangan masyarakat yang mendukung kebijakan efisiensi anggaran 2025. Masyarakat menunjukkan harapan terhadap kebijakan ini sebagai langkah yang dapat membawa dampak positif, meskipun masih terdapat beberapa tantangan yang perlu dihadapi. Dari data yang ada, penulis mengidentifikasi masalah yang ada pada setiap ujaran masyarakat yang setuju akan kebijakan tersebut sebagai berikut:

1. Penyediaan informasi yang jelas untuk meminimalisir kekhawatiran publik

Beberapa masyarakat mengungkapkan bahwa kekhawatiran mengenai kebijakan efisiensi anggaran seringkali muncul karena kurangnya informasi

yang jelas. Sebagai contoh, salah satu tweet menyatakan, "Btw anggaran efisiensi itu ga sama dengan anggaran CPNS ya. Selama belum ada info dari menpanRB/ pres kurang kurangilah fear mongering kayak gini." Ini menunjukkan bahwa dengan adanya penjelasan yang tepat, ketakutan yang tidak berdasar bisa diminimalisir.

2. Penyesuaian Anggaran yang Responsif terhadap Dinamika Kondisi

Perubahan kondisi yang ada mengharuskan penyesuaian anggaran agar lebih tetap sasaran. Seperti yang tercermin dalam salah satu *tweet* yang mengatakan, "Lah ya gak gimana-gimana artinya keadaan udah berubah antara saat ini dan ketika besok cpns dibuka. Namanya situasi dan kondisi itu selalu dinamis. Justru bagus kalau emang buka krn berarti anggaran APBN berjalan sesuai rencana dan tidak terpengaruh sama efisiensi dana." Hal ini menunjukkan bahwa anggaran perlu disesuaikan dengan perubahan situasi agar tetap berjalan sesuai dengan rencana yang telah ditetapkan.

3. Efisiensi anggaran sebagai langkah untuk penyusunan ulang prioritas pemerintah

Beberapa pihak melihat kebijakan ini sebagai upaya untuk memastikan bahwa anggaran yang tersedia digunakan secara lebih efektif dan sesuai dengan kebutuhan yang paling mendesak. Sebagai contoh, dalam salah satu tweet disebutkan, "Sejauh ini saya masih menilai positif langkah efisiensi yg dilakukan pemerintah. Selama ini memang banyak anggaran tidak digunakan dengan tepat. Memang akan ribut dulu tapi akan ada keseimbangan baru. Kementerian2 seharusnya mengatur ulang prioritas penggunaan anggaran."

Pernyataan ini mencerminkan pandangan bahwa efisiensi anggaran memberikan kesempatan untuk memperbaiki alokasi dana yang sebelumnya tidak tepat sasaran dan lebih mengutamakan program yang lebih bermanfaat bagi masyarakat. Dengan demikian, kebijakan ini dapat mendorong pemerintah untuk lebih fokus pada prioritas yang mendesak dan berdampak langsung pada kesejahteraan publik.

4. Menjaga kelangsungan program pemerintah melalui efisiensi anggaran

Efisiensi anggaran juga berperan penting dalam memastikan bahwa program-program pemerintah tetap berjalan meskipun terjadi penghematan dana. Beberapa masyarakat menilai bahwa meskipun kebijakan efisiensi anggaran membawa tantangan, namun program yang memiliki dampak langsung terhadap masyarakat tetap dapat dilaksanakan. Sebagai contoh, dalam salah satu tweet disebutkan, "Seburuk-buruknya pemerintahan pak Jokowi setidaknya pemerintahan gak sampai lumpuh demi program ambisus Makan Bergizi Gratis bahkan efisiensi anggaran saat pandemi aja program pemerintahan pusat dan daerah tetap jalan." Pernyataan ini menunjukkan bahwa meskipun terjadi pengurangan anggaran, pemerintah tetap mampu menjalankan program penting, seperti program bantuan sosial dan infrastruktur, yang tetap mendukung kepentingan publik. Dengan demikian, efisiensi anggaran tidak hanya berfokus pada pengurangan dana, tetapi juga pada upaya untuk mempertahankan kelangsungan program-program yang esensial.

Kebijakan efisiensi anggaran 2025 berpotensi membawa perubahan positif dalam pengelolaan keuangan negara dengan memastikan alokasi dana yang lebih tepat sasaran. Dengan mengatur ulang prioritas anggaran, pemerintah dapat fokus pada program yang lebih bermanfaat bagi masyarakat. Meskipun ada tantangan, efisiensi ini dapat menjaga kelangsungan program penting dan meningkatkan efektivitas penggunaan dana untuk mendukung tujuan pembangunan yang lebih optimal.

4.5.2 Tweet Kategori Negatif

Tweet Negatif dalam data yang digunakan pada penelitian mencerminkan pandangan masyarakat yang tidak mendukung kebijakan efisiensi anggaran 2025. Beberapa masyarakat menyuarakan kekhawatiran mengenai kebijakan ini, terutama terkait dengan pengurangan dana, ketidakjelasan alokasi anggaran, dan dampak langsung terhadap kehidupan masyarakat. Dari data yang ada, penulis mengidentifikasi masalah yang ada pada setiap ujaran masyarakat yang tidak setuju dengan kebijakan tersebut sebagai berikut:

1. Kekhawatiran pemangkasan anggaran pada layanan publik.

Banyak masyarakat yang merasa khawatir dengan pengurangan dana yang dilakukan dalam kebijakan efisiensi anggaran, karena mereka merasa ketidakjelasan dalam alokasi dana dapat berdampak negatif terhadap berbagai sektor yang memerlukan anggaran yang cukup. Sebagai contoh, salah satu tweet menyatakan, “Efisiensi anggaran bisa mengurangi dana untuk sektor penting seperti pendidikan dan kesehatan. Gimana nasib anggaran buat rakyat kecil kalau begini?” Hal ini menunjukkan bahwa masyarakat khawatir

anggaran yang dikurangi bisa mengganggu sektor-sektor vital yang menyentuh langsung kebutuhan rakyat.

2. Penghamburan anggaran untuk kepentingan pejabat.

Banyak tweet negatif yang mengkritik pengelolaan anggaran yang dinilai tidak transparan dan justru digunakan untuk kepentingan pejabat atau kelompok tertentu. Salah satu tweet mencatat, "Efisiensi anggaran malah jadi alasan buat pejabat korup. Mereka tetap dapat dana untuk kepentingan pribadi, sedangkan masyarakat menderita." Ini menggambarkan ketidakpercayaan publik terhadap pengelolaan anggaran yang ada dan kecemasan mengenai penyalahgunaan dana.

3. Pengurangan pendapatan bagi Pegawai Negeri Sipil (PNS)

Kebijakan efisiensi anggaran yang berpotensi mengurangi gaji pekerja di Indonesia termasuk Pegawai Negeri Sipil (PNS) menjadi sorotan bagi banyak orang. Masyarakat merasa bahwa pengurangan anggaran bisa berdampak langsung pada kesejahteraan PNS, yang selama ini mengandalkan gaji tetap. Sebuah tweet menyebutkan, "Gaji PNS bisa dipotong hanya demi efisiensi? Sepertinya pemerintah gak memikirkan kesejahteraan para pekerja negeri." Pandangan ini menunjukkan ketidaksetujuan terhadap kebijakan yang dianggap merugikan para pegawai yang bekerja di sektor publik.

4. Kebijakan penghematan yang menyulitkan masyarakat.

Banyak masyarakat yang merasa bahwa kebijakan efisiensi anggaran justru lebih banyak memberikan dampak negatif kepada mereka, terutama dalam hal pelayanan publik atau pengurangan program sosial yang selama ini

mereka manfaatkan. Sebagai contoh, salah satu tweet mengatakan, "Efisiensi anggaran hanya bikin masyarakat kesulitan. Program bantuan sosial dipotong, sementara pejabat tetap hidup enak." Ini menunjukkan bahwa kebijakan penghematan anggaran dianggap menyulitkan masyarakat bawah yang sangat bergantung pada bantuan sosial atau program-program pemerintah.

Dari analisis tweet negatif, kita dapat melihat bahwa masyarakat yang tidak mendukung kebijakan efisiensi anggaran 2025 lebih cenderung fokus pada dampak langsung yang ditimbulkan terhadap kesejahteraan publik, terutama di sektor-sektor yang mereka anggap penting seperti kesehatan, pendidikan, dan program sosial. Kekhawatiran mengenai transparansi dalam pengelolaan anggaran dan potensi penyalahgunaan dana juga menjadi isu utama yang diangkat dalam banyak tweet negatif. Kebijakan ini dianggap sebagai langkah yang berisiko bagi banyak masyarakat yang merasakan dampak langsung dari pengurangan anggaran.

4.6 Integrasi Penelitian dalam Islam

Perbedaan pendapat dalam masyarakat, termasuk dalam menanggapi kebijakan pemerintah, merupakan hal yang wajar. Keberagaman ini membutuhkan pendekatan yang sistematis untuk memahami setiap sikap yang muncul. Dalam konteks ini, klasifikasi sentimen berperan penting untuk menganalisis tanggapan masyarakat. Prinsip tentang pentingnya pengelompokan dan pemahaman perbedaan dalam islam, dijelaskan dalam Surat Al-Hujurat Ayat 13, yang berbunyi:

يَا أَيُّهَا النَّاسُ إِنَّا خَلَقْنَاهُ مِنْ ذَكَرٍ وَأُنْثَىٰ وَجَعَلْنَاهُمْ شُعُوبًا وَقَبَائِلَ لِتَعَارَفُوا إِنَّ أَكْرَمَكُمْ عِنْدَ اللَّهِ أَتْقَاهُمْ إِنَّ اللَّهَ عَلِيمٌ خَبِيرٌ

“Wahai manusia, sesungguhnya Kami telah menciptakan kamu dari seorang laki-laki dan perempuan. Kemudian, Kami menjadikan kamu berbangsa-bangsa dan bersuku-suku agar kamu saling mengenal. Sesungguhnya yang paling mulia di antara kamu di sisi Allah adalah orang yang paling bertakwa. Sesungguhnya Allah Maha Mengetahui lagi Mahateliti.”(Q.S Al-Hujurat:13)

Menurut Tafsir Kementrian Agama RI, Allah menciptakan manusia dari seorang laki-laki dan perempuan, yaitu Adam dan Hawa, yang menunjukkan bahwa derajat kemanusiaan kita sama, tanpa membedakan suku, bangsa, atau ras. Keberagaman ini dimaksudkan agar kita saling mengenal dan membantu, bukan untuk menghina atau bermusuhan. Allah tidak menyukai kesombongan berdasarkan keturunan, kekayaan, atau jabatan; yang paling mulia di sisi-Nya adalah yang paling bertakwa. Oleh karena itu, tingkatkanlah ketakwaan agar menjadi orang yang mulia di sisi Allah, karena Allah Maha Mengetahui segala sesuatu, baik yang tampak maupun yang tersembunyi (Kemenag, 2022) .

Ayat tersebut mengajarkan bahwa Allah SWT menciptakan manusia dengan berbagai latar belakang berbeda suku, bangsa, dan budaya agar kita bisa saling mengenal dan memahami satu sama lain. Begitu pula dalam penelitian ini, klasifikasi sentimen dilakukan untuk memahami ragam pendapat masyarakat secara lebih terstruktur, tanpa memandang siapa yang menyampaikan, tetapi lebih dari isi tanggapannya. Layaknya pesan dalam ayat tersebut, yang menekankan bahwa kemuliaan seseorang bukan ditentukan dari asal-usulnya melainkan dari ketakwaannya, klasifikasi dalam penelitian ini pun berusaha melihat esensi dari setiap pendapat, bukan latar belakang pemberinya. Dengan begitu, proses klasifikasi ini bukan hanya bersifat teknis, tapi juga mencerminkan nilai-nilai islam

dalam menghargai perbedaan dan membangun pemahaman yang lebih baik di tengah masyarakat yang beragam.

Konsep klasifikasi juga dijelaskan dalam surat Ar-Rum:22 yang berbunyi:

وَمِنْ آيَاتِهِ خُلُقُ السَّمُوتِ وَالْأَرْضِ وَاختِلَافُ السِّنِّكُمْ وَالْوَانِكُمْ إِنَّ فِي ذَلِكَ لَآيَاتٍ لِّلْعَلَمِينَ

“Di antara tanda-tanda (kebesaran)-Nya ialah penciptaan langit dan bumi, perbedaan bahasa dan warna kulitmu. Sesungguhnya pada yang demikian itu benar-benar terdapat tanda-tanda (kebesaran Allah) bagi orang-orang yang berilmu.”(Q.S Al-Rum:22)

Menurut Tafsir Wajiz yang ada pada Lajnah Pentashihan Mushaf Al-Qur'an Badan Litbang dan Diklat Kementerian Agama RI Tahun 2016 ayat tersebut menjelaskan bahwa perbedaan bahasa dan warna kulit manusia merupakan bagian dari tanda-tanda kebesaran dan keesaan Allah. Meskipun manusia diciptakan dengan unsur fisik yang sama seperti lidah, gigi, dan bibir mereka mampu berbicara dalam berbagai bahasa yang berbeda. Demikian pula, meski berasal dari sumber yang sama, manusia memiliki warna kulit yang beragam. Semua keberagaman ini menunjukkan bahwa Allah Maha Kuasa dan Maha Esa dalam menciptakan makhluk-Nya dengan begitu detail dan unik. Hanya orang-orang yang berilmu yang mampu menangkap dan memahami makna di balik keberagaman ini sebagai bukti kebesaran-Nya.

Tafsir tersebut menekankan bahwa keberagaman bahasa dan warna kulit manusia adalah bukti kebesaran dan keesaan Allah. Meskipun manusia diciptakan dengan organ yang sama seperti lidah, bibir, dan gigi-Allah menciptakan kemampuan berbahasa dan ciri fisik yang sangat beragam. Ini menunjukkan bahwa perbedaan bukanlah kelemahan, melainkan bagian dari desain ilahi yang memiliki

makna dan tujuan. Dalam konteks penelitian klasifikasi, prinsip ini sangat relevan. Keberagaman data seperti perbedaan gaya bahasa, pendapat, atau ekspresi emosi, tidak boleh dipandang sebagai hambatan, tetapi justru sebagai informasi berharga yang perlu dikenali dan dikelompokkan dengan cermat. Proses klasifikasi dalam penelitian bertujuan untuk memahami pola-pola dalam keberagaman tersebut, sebagaimana ayat ini mengajarkan bahwa hanya orang-orang berilmu yang mampu melihat hikmah di balik perbedaan. Dengan demikian, klasifikasi bukan sekadar teknik pengolahan data, tetapi juga mencerminkan cara berpikir yang menghargai kompleksitas ciptaan Allah dan berusaha memahami maknanya secara mendalam

Masyarakat sering memberikan tanggapan melalui media sosial, yang menjadi platform utama bagi mereka untuk bebas berpendapat mengenai suatu isu atau informasi. Namun, dalam menyampaikan pendapat, manusia tidak luput dari kesalahan, di mana terkadang tutur kata yang digunakan tidak bijak dan tidak memperhatikan etika komunikasi. Selain itu, seringkali masyarakat berpendapat tanpa mengetahui fakta yang sebenarnya, sehingga hanya berdasarkan asumsi semata. Hal ini dapat menimbulkan efek negatif berupa kesalahpahaman. Larangan untuk menyebarkan informasi yang belum jelas kebenarannya telah disampaikan dalam Al-Qur'an, yaitu pada Surat Al-Isra' ayat 36, yang berbunyi

وَلَا تَقْفُ مَا لَيْسَ لَكَ بِهِ عِلْمٌ إِنَّ السَّمْعَ وَالْبَصَرَ وَالْفُؤَادَ كُلُّ أُولَٰئِكَ كَانَ عَنْهُ مَسْئُولًا

“Janganlah engkau mengikuti sesuatu yang tidak kauketahui. Sesungguhnya pendengaran, penglihatan, dan hati nurani, semua itu akan diminta pertanggungjawabannya.” (Q.S Al-Hujurat:6)

Dalam Tafsir As-Sa'di yang disampaikan oleh pakar tafsir Syaikh Abdurrahman bin Nashir as-Sa'di ayat tersebut menjelaskan bahwa setiap tindakan

dan perkataan haruslah dilakukan dengan penuh kehati-hatian, karena semuanya akan dimintakan pertanggungjawaban. Sebagai seorang hamba, kita harus menyadari bahwa setiap yang kita katakan dan perbuat akan dipertanggungjawabkan di hadapan Allah. Oleh karena itu, kita harus menggunakan anggota tubuh hanya untuk ibadah kepada Allah, mengikhlaskan agama hanya untuk-Nya, dan menjauhi segala yang dibenci-Nya.

Dalam konteks ini, klasifikasi tanggapan masyarakat harus dilakukan dengan hati-hati dan teliti, memastikan bahwa setiap opini yang dikumpulkan telah diverifikasi kebenarannya. Hal ini penting untuk menghindari kesalahan interpretasi atau penyebaran informasi yang tidak akurat. Dengan demikian, penelitian ini membantu pembuat kebijakan untuk mengambil keputusan yang lebih adil, objektif, dan responsif terhadap masyarakat, sesuai dengan prinsip kebenaran dan keadilan yang diajarkan dalam Islam.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Pada penelitian ini telah dilakukan beberapa proses mulai dari pengambilan data, *preprocessing*, pembobotan kata menggunakan TF-IDF serta penggunaan metode *Long Short Term Memory*. Hasil penelitian ini menunjukkan bahwa metode LSTM efektif dalam mengklasifikasikan tanggapan masyarakat terkait Kebijakan Efisiensi Anggaran 2025. Dengan menggunakan dataset yang terdiri dari 2222 data tweet, yang telah dibersihkan dan diberi label sebagai sentimen positif dan negatif, model LSTM mampu mencapai tingkat akurasi yang tinggi.

Berdasarkan hasil pengujian pada tiga skenario rasio pembagian data 70:30, 80:20, dan 90:10 menggunakan model *Long Short-Term Memory* (LSTM) yang memberikan performa secara keseluruhan sangat baik. Pada rasio 70:30, akurasi tertinggi mencapai 93,85% dengan *epoch* 10 dan *batch size* 64, namun seiring bertambahnya *epoch*, performa model cenderung menurun. Pada rasio 80:20, akurasi tertinggi mencapai 94,38% pada *epoch* 10 dengan *batch size* 32, menunjukkan bahwa dengan pembagian data yang lebih seimbang, model LSTM dapat mempelajari pola lebih baik dan menghasilkan hasil yang lebih akurat. Untuk rasio 90:10, meskipun jumlah data pelatihan lebih banyak, akurasi tertinggi mencapai 92,38% dengan *epoch* 10 dan *batch size* 32, sedikit lebih rendah dibandingkan dengan rasio 70:30 dan 80:20.

Akurasi tertinggi yang dicapai pada skenario rasio 80:20 dengan *epoch* 10 dan *batch size* 32 sebesar 94,38% menunjukkan bahwa model LSTM bekerja paling optimal dengan data yang lebih seimbang antara pelatihan dan pengujian. Penggunaan *batch size* 32 pada *epoch* 10 memberikan keseimbangan yang baik antara memori yang digunakan dan waktu pelatihan, sehingga model dapat mempelajari pola dalam data dengan efektif tanpa mengalami *overfitting*. Hasil ini menegaskan pentingnya pengaturan *hyperparameter* yang tepat dalam mencapai performa terbaik pada model LSTM, dengan pertimbangan rasio pembagian data yang lebih seimbang.

5.2 Saran

Penulis menyadari adanya keterbatasan dalam penelitian ini dan sebagai tindak lanjut, penulis memberikan beberapa saran untuk penelitian selanjutnya. Saran tersebut diharapkan dapat mengatasi kekurangan yang ada dengan mempertimbangkan temuan-temuan yang telah dicapai dalam penelitian ini.

1. Menyeimbangkan jumlah data kategori guna mengurangi bias dalam model, dengan menerapkan metode *oversampling* atau SMOTE.
2. Menggunakan metode *Word2Vec* atau *GloVe* sebagai alternatif pembobotan kata yang dapat lebih memahami konteks kata dalam kalimat dibandingkan TF-IDF
3. Menggunakan metode yang lebih canggih seperti GRU atau *Transformer* yang lebih efisien dalam menangani data sekuensial panjang agar dapat membandingkan performanya.

5.3 Rekomendasi

Berdasarkan hasil analisis konten, sebagian besar tanggapan masyarakat terhadap kebijakan efisiensi anggaran 2025 cenderung negatif, dengan banyaknya keluhan terkait dampak kebijakan tersebut. Untuk mengatasi permasalahan yang muncul dari tanggapan negatif tersebut, berikut disajikan beberapa rekomendasi yang dapat diterapkan untuk memperbaiki kebijakan ini dan meningkatkan *respons* positif dari masyarakat.

1. Perjelas Alokasi Anggaran untuk Sektor Penting.
2. Peningkatan Pengawasan terhadap Penggunaan Anggaran.
3. Penjaminan Kesejahteraan PNS dalam Kebijakan Efisiensi.
4. Kebijakan Penghematan yang Tidak Memberatkan Masyarakat.
5. Penyuluhan untuk Meningkatkan Pemahaman Masyarakat.
6. Transparansi dalam Pengalokasian Dana.

DAFTAR PUSTAKA

- Alfarizi, M. I., Syafaah, L., & Lestandy, M. (2022). Emotional Text Classification Using TF-IDF (Term Frequency-Inverse Document Frequency) and LSTM (Long Short-Term Memory). *Jurnal Informatika*, 10(2), 225–232. <https://doi.org/10.30595/juita.v10i2.13262>
- Amna, A., S, W., I, G. I. S., Putra, T. A. E., Wahidin, A. J., Syukrilla, W. A., Wardhani, A. K., Heryana, N., Indriyani, T., & Santoso, L. W. (2023). *Data Mining*. PT Global Eksekutif Teknologi. <https://repository.uinjkt.ac.id/dspace/handle/123456789/71093>
- Anggraini, W. P., & Utami, M. S. (2021). Klasifikasi Sentimen Masyarakat Terhadap Kebijakan Kartu Pekerja Di Indonesia. *Faktor Exacta*, 13(4), 255. <https://doi.org/10.30998/faktorexacta.v13i4.7964>
- Ardianti, P. F., Ramadhani, A. D., & Zein, A. W. (2025). Analisis Pembiayaan Sektor Publik oleh Pemerintah dalam Meningkatkan Pelayanan Sosial Publik di Indonesia: Pendekatan Kualitatif. *Jurnal Penelitian Ilmu-Ilmu Sosial*, 2(10), 412–418. <https://doi.org/10.5281/ZENODO.15493681>
- Azrul, A., Irma Purnamasari, A., & Ali, I. (2024). Analisis Sentimen Pengguna Twitter Terhadap Perkembangan Artificial Intelligence Dengan Penerapan Algoritma Long Short-Term Memory (LSTM). *Jurnal Mahasiswa Teknik Informatika*, 8(1), 413–421. <https://doi.org/10.36040/jati.v8i1.8416>
- Badan Pemeriksa Keuangan. (2025). *Efisiensi Belanja dalam Pelaksanaan Anggaran Pendapatan dan Belanja Negara dan Anggaran Pendapatan dan Belanja Daerah Tahun Anggaran 2025*. Peraturan.Bpk.Go.Id. <https://peraturan.bpk.go.id/Details/313401/inpres-no-1-tahun-2025>
- Choliq, A. (2025). *Dukungan ASN Terhadap Efisiensi Anggaran Menjadi Kunci Utama Keberhasilan*. <https://www.djkn.kemenkeu.go.id/kanwil-rsk/baca-artikel/17532/Dukungan-ASN-Terhadap-Efisiensi-Anggaran-Menjadi-Kunci-Utama-Keberhasilan.html>
- Deby, S. A., & Isnain, A. R. (2021). Teks Dan Analisis Sentimen Pada Chat Grup Whatsapp Menggunakan Long Short Term Memory (LSTM). *Jurnal Teknologi Dan Sistem Informasi*, 3(2), 1. <https://doi.org/doi.org/10.33365/jtsi.v2i4.1248>
- Firdaus, A. A. (2025). *Sentiment Budget Efficiency in Indonesia* [Dataset]. <https://www.kaggle.com/datasets/jocelyndumlao/budget-efficiency-in-indonesia>

- Fitrana, L. A., Linawati, S., Herlinawati, N., Sa'adah, R., & Seimahuria, S. (2024). Analisis Sentimen Pengguna Twitter Terhadap Brand Indosat Menggunakan Metode Naïve Bayes Classifier. *Jurnal Mahasiswa Teknik Informatika*, 8(3), 4291–4297. <https://doi.org/10.36040/jati.v8i3.9866>
- Izzadiana, H. S., Napitupulu, H., & Firdaniza, F. (2023). Peramalan Data Univariat Menggunakan Metode Long Short Term Memory. *Jurnal Sistem Informasi Dan Informatika*, 5(2), 29–39. <https://doi.org/10.37278/sisinfo.v5i2.669>
- Jelodar, H., Wang, Y., Orji, R., & Huang, S. (2020). Deep Sentiment Classification and Topic Discovery on Novel Coronavirus or COVID-19 Online Discussions: NLP Using LSTM Recurrent Neural Network Approach. *IEEE Journal of Biomedical and Health Informatics*, 24(10), 2733–2742. <https://doi.org/10.1109/JBHI.2020.3001216>
- Kemenag, Q. (2022). *Tafsir Al-Quran Kementerian Agama RI QS. Al-Hujurat*. <https://quran.kemenag.go.id/quran/per-ayat/surah/49?from=13&to=18>
- Khumaidi, A., & Nirmala, I. A. (2022). *Algoritma Long Short Term Memory dengan Hyperparameter Tuning: Prediksi Penjualan Produk*. Deepublish.
- Krisnha, R. W. (2024). *Klasifikasi Hate Speech dan Abusive pada Tweet di Platform X* (Issue 4(101)) [Universitas Dinamika]. <https://doi.org/10.25633/sgi.2021.04.02>
- Kurniawan, K., Ceasaro, B., & Sucipto, S. (2024). Perbandingan Fungsi Aktivasi untuk Meningkatkan Kinerja Model LSTM dalam Prediksi Ketinggian Air Sungai. *Jurnal Edukasi Dan Penelitian Informatika*, 10(1). <https://doi.org/10.26418/jp.v10i1.72866>
- Lisanthoni, A., Gunawan, E. L., Adhigiadany, C. A., & Prasetya, A. (2024). Penerapan LSTM dalam Analisis Sentimen Berbasis Lexicon untuk Meningkatkan Sistem Pemantauan Citra PLN di Platform Digital. *Seminar Nasional Sains Data*, 4(1), 581–591. <https://doi.org/10.33005/senada.v4i1>
- Mahmuji, F., Muchammad, C., Cahyo, C., & Nisa, U. (2023). Analisis Sentimen Pada Twitter Terhadap Isu Penundaan Pemilu 2024 Dengan Membandingkan Metode Long Short- Term Memory Dan Naïve Bayes Classifier. *Journal of Mathematics & Information Technology*, 1(2). <https://doi.org/10.35718/equiva.v1i2>
- Mao, Y., Liu, Q., & Zhang, Y. (2024). Sentiment analysis methods, applications, and challenges: A systematic literature review. *Journal of King Saud University - Computer and Information Sciences*, 36(4), 102048. <https://doi.org/10.1016/j.jksuci.2024.102048>

- Maulana, A. (2025). *Dampak Efisiensi Anggaran 2025, dari PHK hingga Potong Beasiswa*. <https://tirto.id/dampak-efisiensi-anggaran-2025-dari-phk-hingga-potong-beasiswa-g8lK>
- Mursyid, R., & Indriyanti, A. D. (2024). Perbandingan Akurasi Metode Analisis Sentimen Untuk Evaluasi Opini Pengguna Pada Platform Media Sosial (Studi Kasus: Twitter). *Journal of Informatics and Computer*, 06(02), 371–383. <https://doi.org/10.26740/jinacs.v6n02.p371-383>
- Musfiroh, M., Tholib, A., & Arifin, Z. (2024). Analisis Sentimen Terhadap Ulasan Aplikasi Shopee di Google Play Store Menggunakan Metode TF-IDF dan Long Short-Term Memory (LSTM). *Journal of Electrical Engineering and Computer*, 6(2). <https://doi.org/10.33650/jeecom.v4i2>
- Nurashila, S. S., Hamami, F., & Kusumasari, T. F. (2023). Perbandingan Kinerja Algoritma Recurrent Neural Network (RNN) dan Long Short-Term Memory (LSTM): Studi Kasus Prediksi Kemacetan Lalu Lintas Jaringan PT XYZ. *Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika*, 8(3), 864–877. <https://doi.org/10.29100/jipi.v8i3.3961>
- Prasetyo, H., Crysdian, C., & Santoso, I. B. (2023). *Deteksi Spam Pada Trending Topik Twitter Berbahasa Indonesia Menggunakan Artificial Neural Network dan Stochastic Gradient Descent*. 15(22), 45–51. <https://doi.org/doi.org/10.35891/explorit.v11i2.1793>
- Pravina, A. M., Cholissodin, I., & Adikara, P. P. (2019). Analisis Sentimen Tentang Opini Maskapai Penerbangan pada Dokumen Twitter Menggunakan Algoritme Support Vector Machine (SVM). *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 3(3), 2789–2797. <https://doi.org/10.1061/9780784415313.ch06>
- Purnasiwi, R. G., Kusrini, & Hanafi, M. (2023). Analisis Sentimen Pada Review Produk Skincare Menggunakan Word Embedding dan Metode Long Short-Term Memory (LSTM). *Journal Of Social Science Research*, 3(2), 11433–11448. <https://doi.org/10.31004/innovative.v3i2>
- Purwanti, A. A., Ramadhani, A. T., & Hayati, K. R. (2024). Pengaruh Media Sosial X Terhadap Interaksi Sosial dan Keterlibatan Kewarganegaraan dan Masyarakat. *Jurnal Hukum Dan Kewarganegaraan*, 3(8). <https://doi.org/doi.org/10.3783/causa.v3i8.3291>
- Rahman, A. (2025, February 19). *Analisis Drone Empirit*. Analisis Percakapan Publik: Efisiensi Anggaran. <https://pers.droneempriit.id/analisis-percakapan-publik-efisiensi-anggaran/>
- Sari, H., Leonarde Ginting, G., & Zebua, T. (2021). Penerapan Algoritma Text Mining dan TF-IDF untuk Pengelompokan Topik Skripsi pada Aplikasi

Repository STMIK Budi Darma. *Terapan Informatika Nusantara*, 2(7), 414–432. <https://doi.org/10.47065/tin.v5i12>

Sejati, W., Bist, A. S., & Tambunan, A. (2023). Pengembangan Analisis Sentimen dalam Rekayasa Software Engineering Menggunakan Tinjauan Literatur Sistematis. *Manajemen, Pendidikan Dan Teknologi Informasi*, 2(1), 95–103. <https://doi.org/10.33050/mentari.v2i1.377>

Septiawan, Y. (2023). Perbandingan Akurasi Metode Deteksi Ujaran Kebencian dalam Postingan Twitter Menggunakan Metode SVM dan Decision Trees yang Dioptimalkan dengan Adaboost. *Jurnal Ilmiah Bidang Ilmu Rekayasa*, 17(2). <https://doi.org/10.5281/zenodo.8175025>

Shehzad, F., Rehman, A., Javed, K., Alnowibet, K. A., Babri, H. A., & Rauf, H. T. (2022). Binned Term Count: An Alternative to Term Frequency for Text Categorization. *Mathematics*, 10(21), 1–25. <https://doi.org/10.3390/math10214124>

Wang, J., & Xu, R. (2023). Performance Analysis of Sentiment Classification Based Neural Network. *Applied and Computational Engineering*, 5(1), 513–518. <https://doi.org/10.54254/2755-2721/5/20230633>

Wesley, R., & Gunawan, R. (2024). Literatur Riview: Metode Deep Learning untuk Analisis Teks. *Jurnal Mahasiswa Teknik Informatika*, 8(5), 11020–11023. <https://doi.org/10.36040/jati.v8i5.11780>

Yulita, W. (2021). Analisis Sentimen Terhadap Opini Masyarakat Tentang Vaksin Covid-19 Menggunakan Algoritma Naïve Bayes Classifier. *Jurnal Data Mining Dan Sistem Informasi*, 2(2), 1. <https://doi.org/10.33365/jdmsi.v2i2.1344>

Zahidin, I., Kanata, B., & Akbar, L. A. S. I. (2024). Perkiraan Suhu Menggunakan Algoritma Recurrent Neural Network Long Short Term Memory. *Journal of Computer System and Informatics (JoSYC)*, 6(1), 244–252. <https://doi.org/10.47065/josyc.v6i1.6242>