



Analisis Sentimen Penghapusan Skripsi sebagai Tugas Akhir Mahasiswa Menggunakan Metode Multi-Layer Perceptron

Haerunnisya Makmur¹, Wulandari², Dewi Fatmarani Surianto^{3*}, Muhammad Fajar B⁴

Program Studi Teknik Komputer, Fakultas Teknik, Universitas Negeri Makassar
Jl. Daeng Tata Raya Parang Tambung, Makassar, Indonesia 90224

**email: dewifatmaranis@unm.ac.id*

(Naskah masuk: 26 Februari 2024; direvisi: 10 Oktober 2024; diterima untuk diterbitkan: 15 Oktober 2024)

ABSTRAK – Indonesia memiliki tingkatan pendidikan dimana salah satunya adalah pendidikan sarjana. Terdapat persyaratan yang harus dilakukan untuk mendapatkan gelar sarjana, salah satunya adalah menyelesaikan tugas akhir dalam bentuk skripsi. Nadiem Makarim, Mendikbudristekdikti dalam pidatonya mengumumkan kebijakan baru di bidang pendidikan mengenai ketidakwajiban bagi mahasiswa untuk menyusun skripsi sebagai syarat kelulusan. Berdasarkan hal tersebut, muncul pro dan kontra dari masyarakat terutama dalam hal lingkup akademis, proses analisis sentiment terkait hal tersebut menjadi hal yang dibutuhkan. Penelitian ini bertujuan untuk memetakan opini masyarakat yang terdapat pada media sosial TikTok dan YouTube terkait penghapusan skripsi menggunakan metode Multi-Layer Perceptron. Tahapan yang dilakukan terdiri dari observasi, pengumpulan data, pelabelan, normalisasi data, preprocessing, pembagian partisi data, pembobotan Term Frequency-Inverse Document Frequency, klasifikasi, dan evaluasi. Akurasi yang diperoleh pada tahap skenario preprocessing yaitu 86% dengan skenario case folding dan stemming. Selanjutnya skenario ini digunakan pada pengujian berdasarkan partisi data dimana diperoleh hasil akurasi tertinggi dengan porsi 90% data latih dan 10% data uji. Akurasi yang diperoleh yakni 94%.

Kata Kunci – Skripsi; Analisis Sentimen; Multi-Layer Perceptron; Opini; Klasifikasi.

Sentiment Analysis of the Elimination of Thesis as Student Final Project Using Multi-Layer Perceptron Method

ABSTRACT – Indonesia has levels of education where one of them is undergraduate education. There are requirements that must be done to get a bachelor's degree, one of which is to complete a final project in the form of a thesis. Nadiem Makarim, Minister of Education, Technology and Higher Education in his speech announced a new policy in the field of education regarding the non-obligation for students to prepare a thesis as a requirement for graduation. Based on this, there are pros and cons from the community especially in terms of the academic sphere, the sentiment analysis process related to this is needed. This research aims to map public opinion contained in TikTok and YouTube social media related to the elimination of thesis using the Multi-Layer Perceptron method. The stages carried out consist of observation, data collection, labeling, data normalization, preprocessing, data partitioning, Term Frequency-Inverse Document Frequency weighting, classification, and evaluation. The accuracy obtained at the preprocessing scenario stage is 86% with case folding and stemming scenarios. Furthermore, this scenario is used in testing based on data partitioning where the highest accuracy results are obtained with a portion of 90% training data and 10% test data. The accuracy obtained is 94%.

Keywords – Thesis; Sentiment Analysis; Multi-Layer Perceptron; Opinion; Classification.

1. PENDAHULUAN

Sistem pendidikan Indonesia melibatkan tingkat pendidikan dari dasar hingga tinggi, salah satunya pendidikan sarjana [1]. Menurut Kamus Besar Bahasa Indonesia (KBBI), sarjana mengacu pada seseorang yang mencapai gelar strata satu setelah menyelesaikan program tertentu di perguruan tinggi [2]. Gelar sarjana seringkali menjadi salah satu persyaratan minimum dalam dunia kerja, tetapi untuk memperolehnya, mahasiswa harus menyelesaikan tugas akhir berupa skripsi [3].

Skripsi dapat diartikan sebagai karya tulis ilmiah yang memuat hasil penelitian dari sebuah permasalahan, yang bisa menumbuhkan pemikiran kritis, keterampilan penelitian, dan kemampuan untuk mensintesis, dan menganalisis informasi [4], [5]. Namun, dalam proses penyusunan skripsi, mahasiswa seringkali menghadapi tantangan yang mengakibatkan mereka tidak dapat lulus tepat waktu [6]. Tantangan tersebut dapat berupa kurangnya motivasi dan kemampuan mahasiswa dalam menulis, kurangnya pemahaman tentang proses skripsi, manajemen waktu, dan masalah kesehatan [7], [8], [9], [10].

Skripsi sebagai tugas akhir mahasiswa merupakan topik yang banyak dibahas dalam literatur akademik. Muncul pro dan kontra dari masyarakat terkait pidato Mendikbud Ristek, Nadiem Makarim dalam episode ke-26 terkait "Transformasi Standar Nasional dan Akreditasi Pendidikan Tinggi" pada acara Merdeka Belajar, yang mengumumkan kebijakan baru tentang ketidakwajiban penyusunan skripsi bagi mahasiswa S1 sebagai syarat kelulusan [11]. Hal ini sesuai dengan Peraturan Mendikbudristek Nomor 53 Tahun 2023 tentang Penjaminan Mutu Pendidikan Tinggi [12].

Penghapusan kewajiban skripsi ini telah menimbulkan kekhawatiran, terutama dalam lingkup akademis, terkait potensi erosi standar akademik dan penurunan kualitas kemampuan intelektual serta kemampuan penelitian para lulusan. Kekhawatiran ini muncul karena skripsi dianggap sebagai salah satu indikator penting untuk mengukur kompetensi ilmiah mahasiswa. Jika penghapusan skripsi ini tidak ditinjau secara komprehensif, dikhawatirkan akan berdampak pada kualitas lulusan dan reputasi akademik institusi pendidikan.

Sebelum dampak tersebut terjadi, analisis sentimen dapat menjadi salah satu solusi yang digunakan untuk memahami pandangan publik, termasuk dari kalangan akademisi, mahasiswa, dan masyarakat umum sehingga dapat diperoleh gambaran lebih jelas tentang bagaimana kebijakan ini diterima atau tidaknya oleh masyarakat.

Dari beberapa penelitian sebelumnya, telah dilakukan analisis sentimen menggunakan metode Jaringan Syaraf Tiruan (JST), khususnya *Multi-Layer Perceptron* (MLP). MLP telah digunakan dalam berbagai penelitian untuk mengklasifikasikan sentimen dengan hasil yang cukup baik. Salah satu penelitian yaitu terkait sentimen terhadap presiden Jokowi dimana datanya berupa *tweet* yang dicari menggunakan kata kunci presiden Jokowi dengan akurasi mencapai 93.26% [13].

Penelitian lain juga menggunakan MLP dengan algoritma *backpropagation* untuk menganalisis sentimen terkait calon presiden Indonesia tahun 2019, dengan akurasi sebesar 90.6% [14]. Selain itu, terdapat pula penelitian yang berfokus pada analisis sentimen terhadap opini pengguna BPJS Kesehatan dengan perolehan akurasi sebesar 87.14% menggunakan metode *backpropagation* [15]. Kedua penelitian ini mendemonstrasikan kemampuan MLP dalam menangkap pola sentimen dari opini publik.

Selain itu, beberapa penelitian lain menggunakan algoritma *backpropagation* untuk analisis sentimen yang dilengkapi dengan metode pembobotan kata menggunakan TF-IDF seperti analisis sentimen terhadap penanggulangan bencana banjir di Provinsi Jawa Barat dengan akurasi sebesar 73.83% [16]. Penelitian berikutnya yaitu analisis sentimen terhadap tanggapan di Twitter terhadap layanan PT Tiki JNE dengan akurasi mencapai 80.27% [17].

Pada beberapa penelitian lain, terdapat penggunaan metode yang berbeda. Analisis sentimen terhadap ulasan pengguna yang terdapat pada aplikasi Google Play Store menggunakan *Support Vector Machine* (SVM) dengan akurasi sebesar 81.46% [18]. Selanjutnya, penelitian analisis sentimen terhadap ulasan aplikasi iPusnas pada Google Play Store menggunakan metode SVM dengan nilai akurasi dicapai 94,24% [19].

Penelitian berikutnya menggunakan metode *Naïve Bayes* untuk mengklasifikasikan sentimen terhadap kinerja Dewan Perwakilan Rakyat (DPR) dengan akurasi yang diperoleh mencapai 80% [20]. Terakhir, terdapat penelitian yang berkaitan dengan analisis sentimen mengenai aplikasi TikTok dengan akurasi sebesar 84% menggunakan SVM [21]. Pada penelitian-penelitian di atas, terdapat variasi dalam langkah-langkah *preprocessing* seperti *case folding*, *tokenization*, *stopwords removal*, *stemming*, *normalization*, dan *duplicate removal*.

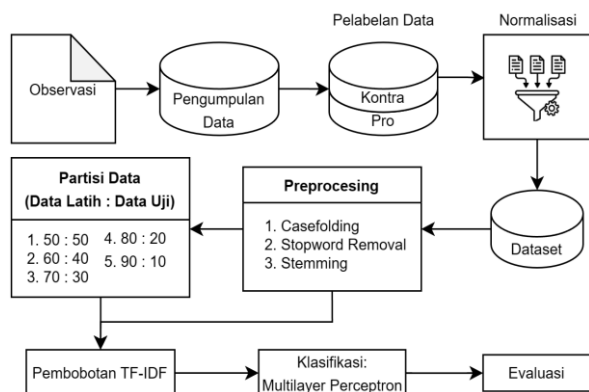
Berdasarkan penelitian-penelitian di atas, dapat disimpulkan bahwa MLP telah banyak digunakan dalam analisis sentimen dengan hasil yang cukup akurat. Termasuk kombinasi MLP dengan algoritma *backpropagation* dan metode pembobotan kata seperti TF-IDF. Namun, kekurangan dalam beberapa penelitian sebelumnya adalah kurangnya optimasi skenario pada tahapan *preprocessing*, dimana

preprocessing dapat memberikan dampak signifikan pada hasil akhir karena proses ini memungkinkan model untuk bekerja dengan teks yang lebih bersih dan representatif.

Oleh karena itu, penelitian ini berfokus pada pengembangan model klasifikasi sentimen menggunakan MLP dengan algoritma *backpropagation*, dengan dua skenario utama. Skenario pertama melibatkan tahapan *preprocessing* data, sementara skenario kedua berkaitan dengan penggunaan partisi data. Penelitian sebelumnya ([14], [15], [20]) masing-masing menggunakan partisi data yang berbeda dalam membangun model, sehingga penelitian ini akan mengeksplorasi penggunaan partisi data untuk mengoptimalkan performa model. Dengan melakukan optimasi yang lebih komprehensif, diharapkan penelitian ini dapat memberikan kontribusi yang signifikan dalam meningkatkan akurasi dan efisiensi model klasifikasi sentimen, terutama dalam konteks analisis sentimen terhadap penghapusan skripsi sebagai tugas akhir mahasiswa.

2. METODE DAN BAHAN

Tahapan penelitian tergambar pada Gambar 1.



Gambar 1. Tahapan penelitian

Observasi

Observasi dilakukan dengan mencari topik terhangat yaitu terkait penghapusan skripsi. Selain itu, dilakukan pula proses wawancara ke beberapa narasumber, baik dosen maupun mahasiswa terkait topik tersebut. Hal ini ditujukan sebagai pengumpulan perspektif yang berbeda terhadap penghapusan skripsi, apakah pro atau kontra. Terdapat 7 dosen dan 12 mahasiswa yang diwawancara yang terdiri dari beberapa kampus berbeda di Indonesia yakni dari Perguruan Tinggi Negeri maupun Perguruan Tinggi Swasta dan dari bidang ilmu berbeda yakni dari Program Studi Manajemen, Informatika, Manajemen Retail, dan Sistem Informasi.

Pengumpulan Data

Pengumpulan data dilakukan secara manual berupa komentar masyarakat dari postingan di media sosial, seperti TikTok dan YouTube yang membahas terkait penghapusan skripsi. Berikut Tabel 1 berisi tautan sumber pengambilan komentar.

Tabel 1. Sumber data

Link	Tanggal
https://www.youtube.com/watch?v=iNg_YP6yit0	30/08/2023
https://www.youtube.com/watch?v=A6451h1QDB0&t=10s	29/08/2023
https://www.tiktok.com/@rianfahardhi/video/7273462913262832901?r=1&t=8h0M9vWZENm	31/08/2023

Pelabelan

Pelabelan dilakukan secara manual oleh tiga orang agar diperoleh hasil yang objektif. Berikut contoh teknik pelabelan dapat dilihat pada Tabel 2.

Tabel 2. Teknik pelabelan

Data	1	2	3	Hasil
sy setuju jika skripsi tidak dihapus	Kontra	Pro	Pro	Pro

Normalisasi Data

Komentar yang diperoleh menggunakan bahasa sehari-hari atau tidak menggunakan bahasa baku yang baik, sehingga memiliki banyak item teks yang unik, salah satu contohnya yaitu emoji. Item tersebut perlu dihilangkan karena tidak mengandung nilai informasi yang penting dalam proses klasifikasi [22]. Selain itu, beberapa kata singkatan juga diubah agar dapat diproses dengan mudah menjadi sebuah informasi.

Data Preprocessing

Preprocessing dilakukan untuk mentransformasikan data yang tidak terstruktur menjadi data yang terstruktur, dengan berbagai metode yakni *case folding*, *stopword removal*, dan *stemming*. Pada *stopword removal* digunakan *Library Natural Language Toolkit* (NLTK) untuk mendapatkan *stopword* bahasa Indonesia. Pada *Stemming* digunakan fungsi *StemmerFactory()* pada *library* Sastrawi.

Pembagian Data Latih dan Data Uji

Data dibagi ke dalam 5 partisi dengan porsi berbeda seperti pada Tabel 3.

Tabel 3. Pembagian partisi data

Data Latih(%)	Data Uji(%)
50	50
60	40
70	30
80	20
90	10

Pembobotan Term-Frequency - Inverse Document Frequency (TF-IDF)

Pembobotan TF-IDF bertujuan untuk konversi data teks menjadi berbentuk numerik [23], dengan rumus perhitungan yang ditunjukkan pada persamaan (1), (2), dan (3).

$$tf_{t,d} = \begin{cases} 1 + \log_{10} tf_{t,d} & \text{if } tf_{t,d} > 0 \\ 0 & \text{otherwise} \end{cases} \quad (1)$$

$$idf_t = \log_{10} \frac{N}{df_t} \quad (2)$$

$$w_{t,d} = tf_{t,d} \times idf_t \quad (3)$$

Keterangan:

$tf_{t,d}$ = frekuensi kemunculan kata t pada dokumen d

N = banyak dokumen

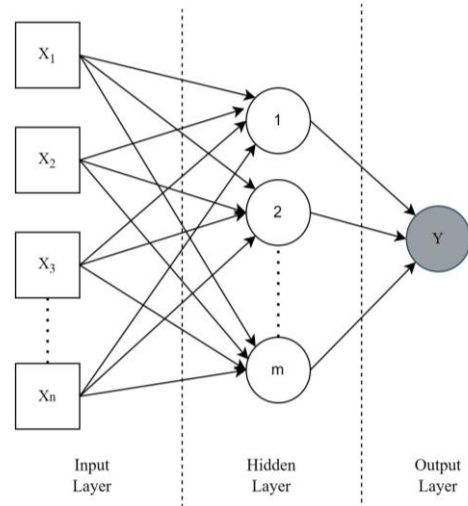
idf_t = banyak dokumen yang memuat t

$w_{t,d}$ = bobot TF.IDF

Pembobotan TF-IDF menggunakan kelas *TfidfVectorizer* dari modul *sklearn feature extraction text* dengan *max features* yaitu 10000. *TfidfVectorizer* digunakan untuk mengonversi koleksi dokumen teks menjadi representasi vektor fitur TF-IDF atau matriks pembobotan. Matriks pembobotan pada data latih harus berukuran sama dengan data uji agar diperoleh dimensi yang sama pada matriks pembobotannya. Sehingga, saat proses pembobotan TF-IDF dilakukan *split* atau membagi *dataset* terlebih dahulu menjadi data latih dan data uji, kemudian matriks pembobotan data latih akan ditransformasi ke data uji.

Klasifikasi

Pada tahap ini dibentuk model yang dapat mengelompokkan data komentar ke dalam dua kategori, baik itu pro maupun kontra. Untuk membentuk model tersebut, digunakan salah satu jenis model MLP yang dapat memodelkan pola dalam data. Model MLP menggunakan algoritma *backpropagation* pada proses pelatihan untuk mengurangi kesalahan antara hasil yang diprediksi oleh jaringan dengan hasil aktual [24]. Arsitektur model MLP pada Gambar 2 menunjukkan bahwa jaringan dibangun dengan jumlah n input, satu lapisan tersembunyi dengan m node, dan 1 node output.



Gambar 2. Arsitektur multi-layer perceptron

Klasifikasi dengan model MLP menggunakan metode *MLPClassifier* dari pustaka *sklearn.neural_network* dalam bahasa pemrograman *Python*. Pembentukan arsitektur dilakukan melalui pengujian parameter *hidden_layer_sizes*, *max_iter*, *activation*, *solver*, dan *learning_rate*. Pengujian *hyperparameter* ini membantu dalam menentukan konfigurasi terbaik untuk model MLP yang sesuai dengan data dan tugas klasifikasi yang spesifik. Hasil pengujian menunjukkan arsitektur MLP yang paling cocok digunakan adalah dengan 250 node pada *hidden layer*, 50 iterasi maksimal, fungsi aktivasi *ReLU*, menggunakan *optimizer Adam*, *learning rate constant*, dan *learning rate init* sebesar 0,001.

Evaluasi

Pengukuran performa atau kinerja model dilakukan berdasarkan data pada *confusion matrix* yang ditampilkan pada Tabel 4, yang selanjutnya diperoleh nilai akurasi, *precision*, *recall*, dan *f1-score* dengan perhitungan menggunakan persamaan (4), (5), (6), dan (7).

Tabel 4. *Confusion matrix* klasifikasi dua kelas

		Prediksi	
		TRUE	FALSE
Aktual	TRUE	True Positive	False Negative
	FALSE	False Positive	True Negative

$$\text{Akurasi} = \frac{TP + TN}{TP + FN + FP + TN} \times 100\% \quad (4)$$

$$\text{Precision} = \frac{TP}{TP + FP} \times 100\% \quad (5)$$

$$\text{Recall} = \frac{TP}{TP + FN} \times 100\% \quad (6)$$

$$\text{F1-score} = \frac{2 \times \text{precision} \times \text{recall}}{\text{precision} + \text{recall}} \times 100\% \quad (7)$$

3. HASIL DAN PEMBAHASAN

Jumlah *dataset* yang dikumpulkan yaitu 600 data komentar dengan rentang waktu unggahan komentar yang diambil yaitu pada tanggal 29 Agustus sampai 5 September 2023. Data teks asli yang diperoleh kemudian dianalisis untuk menentukan label dari setiap data di antara dua kategori label, yaitu “pro” atau “kontra”. Setelah pelabelan awal, data kembali dipilih sehingga diperoleh *dataset* yang terdiri dari 500 komentar dengan jumlah data yang *balance* untuk setiap kategori label, yaitu masing-masing 250 data. Contoh data tergambar pada Tabel 5.

Tabel 5. Contoh data komentar beserta labelnya

No	Teks Data Asli	Label
1	setuju banget sih kalau progam skripsi di hapus dan di ganti demi menyelamatkan dunia pendidikan dari joki2 skripsi	Pro
2	sy setuju jika skripsi tidak dihapus	Kontra

Selanjutnya yaitu tahap normalisasi dapat dilihat pada Tabel 6. Kata *typo* seperti “program” diubah menjadi “program”, kata “di hapus” diubah menjadi “dihapus” sesuai dengan penulisan kata “di” dalam bahasa Indonesia, serta dilakukan penghapusan angka yang membuat kata tidak memiliki nilai seperti “joki2” diubah menjadi “joki”.

Tabel 6. Contoh data hasil normalisasi

No	Teks Hasil Normalisasi
1	setuju banget sih kalau program skripsi dihapus dan diganti demi menyelamatkan dunia pendidikan dari joki skripsi
2	saya setuju jika skripsi tidak dihapus

Data ternormalisasi ini yang selanjutnya akan digunakan pada tahap *preprocessing*. *Preprocessing* data dilakukan dengan 5 skenario berbeda pada masing-masing tahapnya seperti yang ditunjukkan pada Tabel 7.

Tabel 7. Contoh data hasil *preprocessing*

No	Tahap <i>Preprocessing</i>	Hasil
1	Tanpa <i>Preprocessing</i>	Setuju. Anak aku tamat kuliah 8 tahun karena skripsi tidak kunjung lulus. Buat aku juga turun tangan menghubungi dosen. Kalau tidak anak aku bisa di DO. kasihan banget tuh anak
2	<i>Case folding</i>	setuju. anak aku tamat kuliah 8 tahun karena skripsi tidak kunjung lulus. buat aku juga turun tangan menghubungi

No	Tahap <i>Preprocessing</i>	Hasil
3	<i>Case folding</i> → <i>Stopword Removal</i>	dosen. kalau tidak anak aku bisa di do. kasihan banget tuh anak setuju. anak tamat kuliah 8 skripsi kunjung lulus. turun tangan menghubungi dosen. anak do. kasihan banget tuh anak
4	<i>Case folding</i> → <i>Stemming</i>	tuju anak aku tamat kuliah 8 tahun karena skripsi tidak kunjung lulus buat aku juga turun tangan hubung dosen kalau tidak anak aku bisa di do kasihan banget tuh anak
5	<i>Case folding</i> → <i>Stopword Removal</i> → <i>Stemming</i>	tuju anak tamat kuliah 8 skripsi kunjung lulus turun tangan hubung dosen anak do kasihan banget tuh anak

Skenario pertama yaitu tanpa tahap *preprocessing*. Skenario kedua hanya menerapkan *case folding*. Pada skenario ketiga, dilakukan *stopword removal* setelah *case folding* dimana hasil yang didapatkan terdapat beberapa kata yang dihilangkan, seperti kata “aku”, “tidak”, “tahun”, “buat”, dan “kalau”. Dilakukan *case folding* terlebih dahulu agar kata yang dihapus konsisten. Seperti misalnya jika pada daftar *stopword* berisi kata “tidak”, sedangkan kata pada data adalah “Tidak” maka kata “Tidak” tersebut tidak dihilangkan karena dianggap kata yang berbeda.

Skenario keempat dilakukan *stemming* setelah *case folding*. Dari hasil yang diperoleh, kata “setuju” berubah ke bentuk kata dasarnya yaitu “tuju”. Adapun pada skenario kelima dilakukan dua proses setelah *case folding* yaitu *stopword removal* lalu *stemming*. Hasil yang diperoleh seperti pada skenario kedua, tetapi kata imbuhan diubah ke bentuk kata dasarnya, seperti “menghubungi” berubah menjadi “hubung”.

Setelah melakukan tahap *preprocessing*, dilakukan tahap ekstraksi fitur dengan menggunakan pembobotan TF-IDF. Proses TF-IDF dilakukan dengan memecah data menjadi daftar kata yang kemudian disajikan dalam bentuk matriks pembobotan. Matriks pembobotan berisi bobot penilaian kata pada data latih. Jumlah kata dalam matriks pembobotan bergantung pada jumlah kata dalam data latih. Selain itu, semakin tinggi bobot penilaian kata pada matriks pembobotan menunjukkan bahwa kata tersebut cukup unik atau spesifik pada dokumen tertentu.

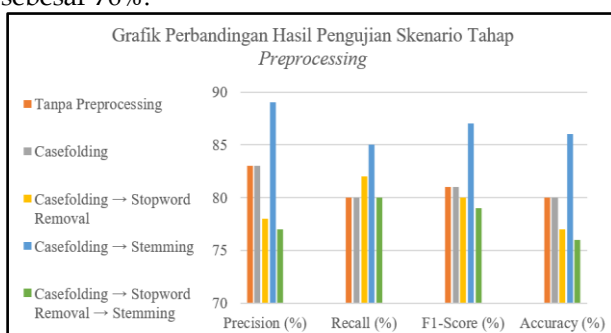
Selanjutnya yaitu proses klasifikasi dengan menggunakan *MLPClassifier*. Pada proses klasifikasi yang pertama yaitu berdasarkan skenario tahap *preprocessing* untuk menentukan skenario *preprocessing* yang memberikan hasil paling

maksimal. Klasifikasi dilakukan dengan membagi data menjadi 80% data latih dan 20% data uji karena merupakan praktik umum dalam pengujian model mengacu pada penelitian [19] yang menggunakan partisi data yang sama. Pada penelitian ini pula dilakukan skenario pengujian porsi data latih dan data uji yang dijelaskan lebih lanjut pada Tabel 9. Hasil pengujian berdasarkan skenario tahap *preprocessing* dapat dilihat pada Tabel 8.

Tabel 8. Hasil pengujian berdasarkan skenario pembagian tahap *preprocessing*

No	Skenario	Precision (%)	Recall (%)	F1-Score (%)	Akurasi (%)
1	Tanpa <i>Preprocessing</i>	83	80	81	80
2	<i>Case folding</i>	83	80	81	80
3	<i>Case folding</i> → <i>Stopword Removal</i>	78	82	80	77
4	<i>Case folding</i> → <i>Stemming</i>	89	85	87	86
5	<i>Case folding</i> → <i>Stopword Removal</i> → <i>Stemming</i>	77	80	79	76

Berdasarkan Tabel 8, nilai akurasi, *precision*, *recall*, dan *f1-score* pada skenario pertama dan kedua bernilai sama. Hal ini terjadi karena seluruh kata pada *dataset* tanpa *preprocessing* yang masih mengandung huruf kapital diubah menjadi huruf kecil pada saat proses TF-IDF sehingga tidak memiliki hasil yang berbeda dengan *dataset* yang melalui tahap *case folding* terlebih dahulu. Adapun akurasi pada skenario ketiga yaitu 77%, skenario keempat yaitu 86%, dan skenario kelima yaitu sebesar 76%.



Gambar 3. Grafik perbandingan hasil pengujian berdasarkan skenario tahap *preprocessing*

Berdasarkan grafik pada Gambar 3, skenario keempat yaitu *case folding* dan *stemming* memiliki tingkat *precision* 89%, *recall* 85%, *f1-score* 87%, dan akurasi 86% tertinggi dibandingkan dengan skenario lain yang menggunakan tahap *stopword removal*.

Penelitian sebelumnya oleh Gunawan dkk [16] yang juga menggunakan MLP dengan tahap *stopword*

removal menghasilkan akurasi sebesar 73,83% yang masih tergolong rendah. Dalam penelitian ini, berdasarkan skenario yang dilakukan, terbukti bahwa penggunaan *stopword removal* justru mengurangi akurasi model. Hal ini menunjukkan bahwa hasil akurasi sistem belum dapat diperoleh secara maksimal dengan pendekatan tersebut.

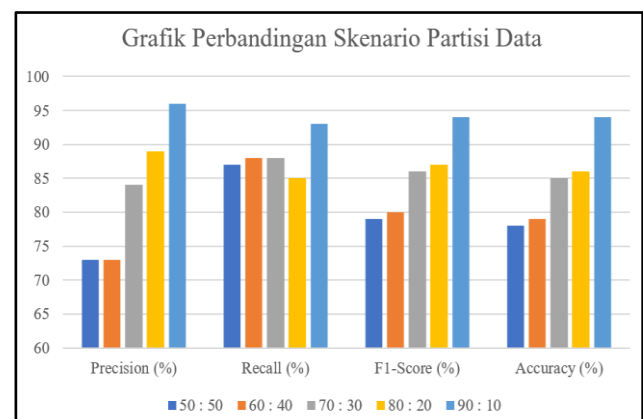
Hal ini disebabkan karena pada proses *stopword removal* dilakukan penghapusan kata yang dapat mempengaruhi kinerja model karena memiliki informasi penting dalam menentukan kategori data seperti penelitian oleh [25] yang juga menunjukkan penurunan akurasi ketika menggunakan *stopword removal*. Hasil yang diperoleh pada Tabel 7, kalimat berupa "anak aku tamat kuliah 8 tahun karena skripsi tidak kunjung lulus" berubah menjadi "anak tamat kuliah 8 skripsi kunjung lulus" setelah melalui proses *stopword removal*. Kedua kalimat tersebut memiliki arti yang berbeda. Kata "tidak" yang dihapus membuat kalimat yang melalui proses *stopword removal* memiliki makna yang sangat berbeda dibandingkan dengan kalimat aslinya.

Berdasarkan hasil yang telah diperoleh, maka untuk pengujian klasifikasi selanjutnya akan digunakan *dataset* yang telah melalui tahap *case folding* dan *stemming*.

Pada proses pengujian klasifikasi yang kedua yaitu berdasarkan pembagian jumlah data latih dan data uji ke dalam 5 partisi data dengan proporsi yang berbeda. Berikut hasil pengujian pada Tabel 9.

Tabel 9. Hasil pengujian berdasarkan skenario partisi data

No	Partisi Data (%)		Precision (%)	Recall (%)	F1-Score (%)	Accuracy (%)
	Latih	Uji				
1	50	50	73	87	79	78
2	60	40	73	88	80	79
3	70	30	84	88	86	85
4	80	20	89	85	87	86
5	90	10	96	93	94	94



Gambar 4. Grafik perbandingan hasil pengujian berdasarkan skenario partisi data

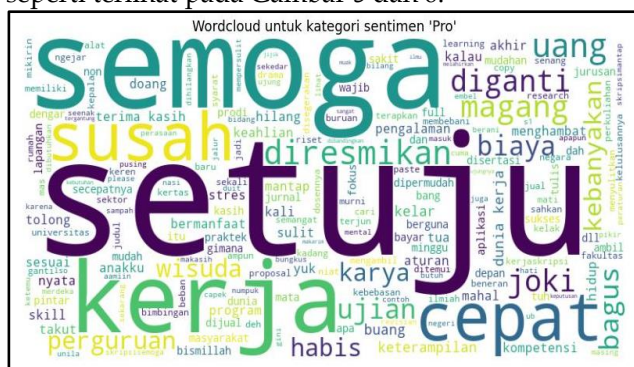
Berdasarkan Tabel 9 dan grafik pada Gambar 4, skenario pertama memiliki akurasi terendah yaitu 78% dan skenario kelima memiliki akurasi tertinggi yaitu 94%. Nilai akurasi yang diperoleh semakin meningkat jika jumlah data latih yang digunakan lebih banyak. Hal ini terjadi ketika data latih yang digunakan semakin banyak maka kinerja model juga semakin meningkat karena lebih banyak mendapatkan informasi saat dilatih [26]. Penelitian yang dilakukan oleh [27] dan [28] juga menunjukkan bahwa akurasi tertinggi dicapai pada percobaan dengan partisi data 90:10. Sehingga, penggunaan 90% data latih dan 10% data uji memiliki nilai akurasi tertinggi, begitupun nilai *precision* 96%, *recall* 93%, dan *f1-score* 94% yang lebih tinggi dibandingkan dengan skenario partisi data yang lain. Berikut *confusion matriks* hasil pengujian yang ditampilkan pada Tabel 10.

Tabel 10. *Confusion matrix* pengujian partisi data 90:10

		Prediksi	
		Kontra	Pro
Aktual	Kontra	2	25
	Pro	1	22

Berdasarkan Tabel 10, dapat dilihat bahwa terdapat 50 data uji yang terdiri dari 23 data dengan label pro dan 27 lainnya dilabeli kontra. Hasil yang diperoleh menunjukkan bahwa jumlah yang diprediksi benar sebagai pro adalah 22, dan 1 data lainnya diprediksi kontra. Adapun pada data yang dilabeli kontra, sebanyak 25 yang diprediksi benar sebagai kontra dan 2 data lainnya salah diprediksi sebagai pro.

Untuk menganalisa kata-kata yang sering muncul pada kategori data pro maupun kontra, digunakan representasi visual *wordcloud*. Representasi *wordcloud* dibuat dari *dataset* hasil *preprocessing* tahap *case folding* dan *stopword removal* dengan menghilangkan kata-kata umum terkait topik yang dianggap tidak menggambarkan sentimen seperti kata "skripsi", "mahasiswa", "kampus", "dosen", "menteri" dan sebagainya. *Wordcloud* untuk dua kategori sentimen seperti terlihat pada Gambar 5 dan 6.



Gambar 5. *Wordcloud* kategori pro

Pada Gambar 5, sentimen kategori pro terhadap penghapusan skripsi menampilkan kata kunci yang sering muncul yaitu “setuju”, “semoga”, “diresmikan”, “bagus” dan lain sebagainya. Kata-kata tersebut mencerminkan bahwa masyarakat antusias menyetujui kebijakan tersebut. Selain itu, terdapat kata “cepat”, “wisuda”, “kerja”, “susah”, dan “magang” yang menggambarkan persepsi opini masyarakat bahwa dengan dihapusnya skripsi, mahasiswa bisa cepat wisuda dan tidak susah mendapatkan kerja karena mendapat pengalaman melalui magang sebagai pengganti skripsi.

Sementara itu, *wordcloud* pada Gambar 6 yang mempresentasikan opini kontra masyarakat terhadap penghapusan skripsi menunjukkan kata-kata seperti “sulit”, “susah”, “berpikir”, “gimana” dan sebagainya. Kata-kata ini mencerminkan kekhawatiran bahwa pengganti skripsi sebagai tugas akhir lebih susah dan sulit.



Gambar 6. *Wordcloud* kategori kontra

Penelitian ini memiliki kontribusi baik secara praktis maupun teoritis. Secara praktis, hasil penelitian ini dapat dijadikan acuan bagi para pemangku kebijakan, khususnya pemerintah, dalam mengambil keputusan yang berkaitan dengan isu penghapusan skripsi karena memberikan gambaran tentang opini publik mengenai kebijakan tersebut. Dari sisi teoritis, penelitian ini berkontribusi melalui penggunaan metode MLP dengan algoritma *backpropagation*, yang dieksplorasi melalui dua skenario berbeda, yakni penerapan tahap *preprocessing* serta variasi partisi data latih dan data uji. Pendekatan ini memungkinkan evaluasi yang lebih mendalam terhadap efektivitas metode tersebut dalam konteks penelitian.

4. KESIMPULAN

Dari hasil studi yang telah dilakukan dengan menerapkan metode MLP dalam klasifikasi analisis sentimen opini masyarakat terhadap penghapusan skripsi sebagai tugas akhir, diperoleh hasil yang menunjukkan tingkat akurasi yang baik. Pada tahap *preprocessing* diperoleh hasil tertinggi dengan nilai *precision* 89%, *recall* 85%, *f1-score* 87%, dan akurasi 86% pada skenario *case folding* dan *stemming*.

Selanjutnya skenario ini digunakan pada pengujian berdasarkan partisi data dengan hasil akurasi tertinggi menggunakan 90% data latih dan 10% data uji. Akurasi yang diperoleh sebesar 94% dengan nilai *precision* 96%, *recall* 93%, dan *f1-score* 94%.

Keterbatasan penelitian ini terletak pada proses pelabelan yang dilakukan secara manual oleh tiga orang. Selain itu, model yang digunakan tidak menjalani optimalisasi parameter. Sebagai rekomendasi, pelabelan data perlu melibatkan lebih banyak pihak untuk meningkatkan validitas, serta disarankan melakukan skenario percobaan dengan optimalisasi *hyperparameter* pada model MLP untuk meningkatkan akurasi dan performa model.

DAFTAR PUSTAKA

- [1] E. B. Prihastari, S. Sukestiyarno, and K. Kartono, "Kajian Literasi Statistik pada Jenjang Pendidikan di Indonesia," *mendidik*, vol. 8, no. 2, pp. 290–299, Oct. 2022.
- [2] "Arti kata sarjana - Kamus Besar Bahasa Indonesia (KBBI) Online." Accessed: Dec. 03, 2023. [Online]. Available: <https://kbbi.web.id/sarjana>
- [3] A. E. R. Fadillah, "Stres dan Motivasi Belajar pada Mahasiswa Psikologi Universitas Mulawarman yang Sedang Menyusun Skripsi," *Psikoborneo*, vol. 1, no. 3, 2013.
- [4] H. P. Tiwari, "Writing Thesis in English Education: Challenges Faced by Students," *Jnl. NELTA Gandaki*, vol. 1, pp. 45–52, Feb. 2019.
- [5] I. W. Djatmiko, *Strategi Penulisan Skripsi, Tesis, Disertasi Bidang Pendidikan*, 1st ed. Yogyakarta: UNY Press, 2018.
- [6] S. Suhandiah, A. Ayuningtyas, and P. Sudarmaningtyas, "Tugas Akhir dan Faktor Stres Mahasiswa," *Jurnal Analisis Sistem Pendidikan Tinggi Indonesia (JAS-PT)*, vol. 5, no. 1, Art. no. 1, Jul. 2021.
- [7] R. Untari, N. Alawiyah, L. Permatasari, F. Sulistiyarini, and S. Q. Melati, "Faktor-Faktor Penghambat Mahasiswa Dalam Menyusun Skripsi," *Academica: Journal of Multidisciplinary Studies*, vol. 6, no. 2, pp. 189–204, Dec. 2022.
- [8] C. L. Coruth, L. D. Boyd, J. N. August, and A. N. Smith, "Perceptions of Dental Hygienists About Thesis Completion in Graduate Education," *Journal of Dental Education*, vol. 83, no. 12, Art. no. 12, Dec. 2019.
- [9] C. Bazan and S. Bayona-Oré, "Why Students Find It Difficult to Finish their Theses?," *Journal of Turkish Science Education*, vol. 17, no. 4, Dec. 2020.
- [10] K. Gamage, S. Hussain, and Q. H. Abbasi, "Final Year Project Allocations for Undergraduate Engineering Students In TNE Programs," in *Proceedings of the 6th Teaching & Education Conference, Vienna, International Institute of Social and Economic Sciences*, Oct. 2018.
- [11] Valorant, *Merdeka Belajar eps 26: Transformasi Standar Nasional dan Akreditasi Pendidikan Tinggi*, (2023). Accessed: Oct. 03, 2023. [Online Video]. Available: <https://www.youtube.com/watch?v=vNTmG4OIQZ4>
- [12] Republik Indonesia, *Peraturan Menteri Pendidikan, Kebudayaan, Riset, dan Teknologi Republik Indonesia Nomor 53 Tahun 2023 Tentang Penjaminan Mutu Pendidikan Tinggi*. 2023, pp. 1–45.
- [13] N. Munasatya and S. Novianto, "Natural Language Processing untuk Analisis Sentimen Presiden Jokowi Menggunakan Multi Layer Perceptron," *Techno.com*, vol. 19, no. 3, pp. 237–244, Aug. 2020.
- [14] I. F. Rozi, Y. Pramitarini, and N. Puspitasari, "Analisis Mengenai Calon Presiden Indonesia 2019 di Twitter Menggunakan Metode Backpropagation," *JIP*, vol. 6, no. 2, Art. no. 2, Feb. 2020.
- [15] R. B. P. Ardika, B. Irawan, and C. Setianingsih, "Analisis Sentimen Data Pada BPJS Kesehatan Dengan Metode Backpropagation Neural Network," *eProceedings of Engineering*, vol. 7, no. 2, Art. no. 2, 2020.
- [16] A. L. S. A.-Z. Gunawan, Jondri, and K. M. Lhaksaman, "Analisis Sentimen pada Media Sosial Twitter terhadap Penanganan Bencana Banjir di Jawa Barat dengan Metode Jaringan Saraf Tiruan," *eProceedings of Engineering*, vol. 8, no. 2, Art. no. 2, Apr. 2021.
- [17] S. F. Aliyah, H. Yasin, S. Suparti, B. Warsito, and T. Widiharhi, "Analisis Sentimen Pt Tiki Jalur Nugraha Ekakurir (PT Tiki JNE) pada Media Sosial Twitter Menggunakan Model Feed Forward Neural Network," *Jurnal Statistika Universitas Muhammadiyah Semarang*, vol. 8, no. 2, Art. no. 2, Nov. 2020.
- [18] L. B. Ilmawan and M. A. Mude, "Perbandingan Metode Klasifikasi Support Vector Machine dan Naïve Bayes untuk Analisis Sentimen pada Ulasan Tekstual di Google Play Store," *ILKOM Jurnal Ilmiah*, vol. 12, no. 2, pp. 154–161, 2020.
- [19] F. S. Lestari, H. Harliana, M. M. Huda, and T. Prabowo, "Sentiment Analysis of iPusnas Application Reviews on Google Play Using Support Vector Machine," *Proceedings of the International Seminar on Business, Education and Science*, vol. 1, pp. 178–188, Oct. 2022.
- [20] D. Duei Putri, G. F. Nama, and W. E. Sulistiono, "Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (DPR) Pada Twitter Menggunakan Metode Naive Bayes Classifier," *JITET*, vol. 10, no. 1, Art. no. 1, Jan. 2022.

- [21] F. A. Indriyani, A. Fauzi, and S. Faisal, "Analisis Sentimen Aplikasi Tiktok Menggunakan Algoritma Naïve Bayes dan Support Vector Machine," *TEKNOSAINS: Jurnal Sains, Teknologi dan Informatika*, vol. 10, no. 2, Art. no. 2, Jul. 2023.
- [22] N. Fitriyah, B. Warsito, and D. A. I. Maruddani, "Analisis Sentimen Gojek pada Media Sosial Twitter dengan Klasifikasi Support Vector Machine (SVM)," *Jurnal Gaussian*, vol. 9, no. 3, pp. 376–390, Aug. 2020.
- [23] S. U. Masruroh, *Fast Text Dan Word2vec Pada Query Kesamaan Semantik Sistem Temu Kembali Informasi*. Deepublish, 2022.
- [24] E. Erdem and F. Bozkurt, "A Comparison of Various Supervised Machine Learning Techniques for Prostate Cancer Prediction," *European Journal of Science and Technology*, no. 21, Art. no. 21, Jan. 2021.
- [25] S. J. Angelina, A. B. P. Negara, and H. Muhardi, "Analisis Pengaruh Penerapan Stopword Removal Pada Performa Klasifikasi Sentimen Tweet Bahasa Indonesia," *Jurnal Aplikasi dan Riset Informatika*, vol. 2, no. 1, pp. 165–173, 2023.
- [26] M. E. Koibur, A. W. Murdiyanto, Z. Munawar, G. P. Insany, H. E. Manurung, D. Karmana, H. Baturohmah, F. Reba, R. W. Simbolon, R. Asrianto, S. A. Mandowen. *Sains Data: Strategi, Teknik, dan Model Analisis Data*. Kaizen Media Publishing, 2023.
- [27] F. B. Saputra, M. Kallista, and C. Setianingsih, "Deteksi Social Distancing Dan Penggunaan Masker Di Restoran Menggunakan Algoritma Residual Network (RESNET)," *eProceedings of Engineering*, vol. 10, no. 1, pp. 284–295, 2023.
- [28] B. N. Azmi, A. Hermawan, and D. Avianto, "Analisis Pengaruh Komposisi Data Training dan Data Testing pada Penggunaan PCA dan Algoritma Decision Tree untuk Klasifikasi Penderita Penyakit Liver," *JTIM: Jurnal Teknologi Informasi dan Multimedia*, vol. 4, no. 4, Art. no. 4, Feb. 2023.