

**IMPLEMENTASI *FEW-SHOT LEARNING* UNTUK PREDIKSI KALIMAT
SOLUSI DARI MASALAH PADA ARTIKEL ILMIAH MENGGUNAKAN
MODEL *LARGE LANGUAGE MODELS* (LLM)**

LAPORAN TUGAS AKHIR

Laporan ini disusun guna memenuhi salah satu syarat untuk menyelesaikan
program studi Teknik Informatika S-1 pada Fakultas Teknologi Industri
Universitas Islam Sultan Agung



Disusun Oleh:

EKA NURUL AZIZAH

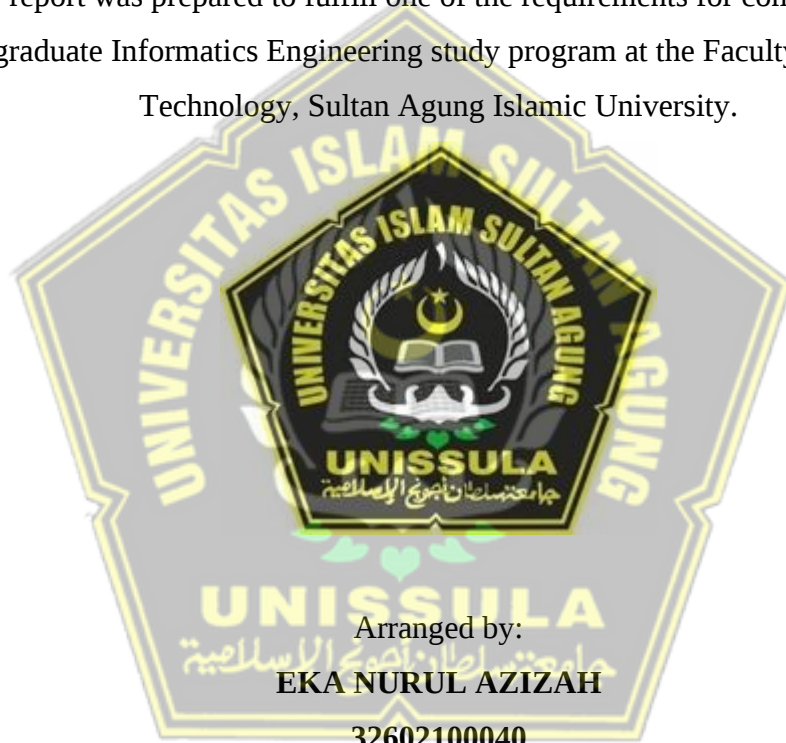
32602100040

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS TEKNOLOGI INDUSTRI
UNIVERSITAS ISLAM SULTAN AGUNG
SEMARANG
2025**

**IMPLEMENTATION OF FEW-SHOT LEARNING TO PREDICTION OF
SENTENCES FOR SOLUTIONS TO PROBLEMS IN SCIENTIFIC
ARTICLES USING LARGE LANGUAGE MODELS (LLM)**

FINAL PROJECT

This report was prepared to fulfill one of the requirements for completing the Undergraduate Informatics Engineering study program at the Faculty of Industrial Technology, Sultan Agung Islamic University.



Arranged by:
EKA NURUL AZIZAH

32602100040

**MAJORING OF INFORMATICS ENGINEERING
FACULTY OF INDUSTRIAL TECHNOLOGY
SULTAN AGUNG ISLAMIC UNIVERSITY**

2025

**LEMBAR PENGESAHAN
TUGAS AKHIR**

**IMPLEMENTASI *FEW-SHOT LEARNING* UNTUK PREDIKSI KALIMAT
SOLUSI DARI MASALAH PADA ARTIKEL ILMIAH MENGGUNAKAN
MODEL *LARGE LANGUAGE MODELS* (LLM)**

**Eka Nurul Azizah
NIM 32602100040**

Telah dipertahankan di depan tim penguji ujian sarjana tugas akhir
Program Studi Teknik Informatika
Universitas Islam Sultan Agung
Pada tanggal : 25 Februari

TIM PENGUJI UJIAN SARJANA :

Ir. Sri Mulyono M.Eng

NIDN. 0626066601

(Ketua Penguji)

Dedy Kurniadi, M.Kom

NIDN. 0622058802

(Anggota Penguji)

Sam Farisa Chaerul Haviana, ST, M.Kom

NIDN. 0628028602

(Pembimbing)

11-03-2025

11-03-2025

11-03-2025

Semarang, 5 Maret 2025

Mengetahui,

Kaprodik Teknik Informatika
Universitas Islam Sultan Agung

Moch. Taufik, S.T., M.IT

NIDN. 0622037302



SURAT PERNYATAAN KEASLIAN TUGAS AKHIR

Yang bertanda tangan dibawah ini :

Nama : Eka Nurul Azizah

NIM : 32602100040

Judul Tugas Akhir : IMPLEMENTASI *FEW-SHOT LEARNING* UNTUK
PREDIKSI KALIMAT SOLUSI DARI KALIMAT
MASALAH PADA ARTIKEL ILMIAH
MENGUNAKAN *LARGE LANGUAGE MODELS*
(LLM)

Dengan baliwa ini saya menyatakan bahwa judul dan isi Tugas Akhir yang saya buat dalam rangka menyelesaikan Pendidikan Strata Satu (S1) Teknik Informatika tersebut adalah asli dan belum pernah diangkat, ditulis ataupun dipublikasikan oleh siapapun baik keseluruhan maupun sebagian, kecuali yang secara tertulis diacu dalam naskah ini dan disebutkan dalam daftar pustaka, dan apabila di kemudian hari ternyata terbukti bahwa judul Tugas Akhir tersebut pernah diangkat, ditulis ataupun dipublikasikan, maka saya bersedia dikenakan sanksi akademis. Demikian surat pernyataan ini saya buat dengan sadar dan penuh tanggung jawab.

UNISSULA
جامعة سلطان أبوبنح الإسلامية

Semarang, 13 Maret 2025

Yang Menyatakan



Eka Nurul Azizah

PERNYATAAN PERSETUJUAN PUBLIKASI KARYA ILMIAH

Saya yang bertanda tangan dibawah ini :

Nama : Eka Nurul Azizah

NIM : 32602100040

Program Studi : Teknik Informatika

Fakultas : Teknologi industri

Dengan ini menyatakan Karya Ilmiah berupa Tugas akhir dengan Judul : IMPLEMENTASI *FEW-SHOT LEARNING* UNTUK PREDIKSI KALIMAT SOLUSI DARI MASALAH PADA ARTIKEL ILMIAH MENGGUNAKAN MODEL *LARGE LANGUAGE MODELS (LLM)*

Menyetujui menjadi hak milik Universitas Islam Sultan Agung serta memberikan Hak bebas Royalti Non-Eksklusif untuk disimpan, dialihmediakan, dikelola dan pangkalan data dan dipublikasikan diinternet dan media lain untuk kepentingan akademis selama tetap menyantumkan nama penulis sebagai pemilik hak cipta. Pernyataan ini saya buat dengan sungguh-sungguh. Apabila dikemudian hari terbukti ada pelanggaran Hak Cipta/Plagiarisme dalam karya ilmiah ini, maka segala bentuk tuntutan hukum yang timbul akan saya tanggung secara pribadi tanpa melibatkan Universitas Islam Sultan agung.

Semarang, 13 Maret 2025

Yang menyatakan,



Eka Nurul Azizah

KATA PENGANTAR

Segala puji dan syukur saya panjatkan ke hadirat Allah SWT atas limpahan rahmat dan karunia-Nya, sehingga saya dapat menyelesaikan Tugas Akhir ini dengan judul “Implementasi *Few-shot Learning* untuk prediksi kalimat solusi dari kalimat masalah pada artikel ilmiah menggunakan LLM”. Tugas Akhir ini disusun sebagai salah satu syarat kelulusan dalam menempuh studi serta untuk memperoleh gelar sarjana (S-1) pada Program Studi Teknik Informatika, Fakultas teknologi Industri, Universitas Islam Sultan Agung Semarang.

Dalam penyusunan Tugas Akhir ini, saya mendapatkan banyak bantuan, baik dalam aspek materi maupun teknis dari berbagai pihak. Oleh karena itu, dengan segala kerendahan hati, saya ingin mengucapkan terima kasih kepada :

1. Rektor UNISSULA Bapak Prof. Dr. H. Gunarto, SH., M.Hum yang mengizinkan penulis menimba ilmu di kampus ini
2. Dekan Fakultas Teknologi Industri Ibu Dr. Ir. Hj. Novi Marlyana, S.T., M.T., IPU., ASEAN.Eng
3. Dosen Pembimbing Bapak Sam Farisa Chaerul Haviana, ST, M.Kom yang telah meluangkan waktu, membimbing, dan memberi ilmu.
4. Orang tua penulis Bapak Romadlon dan Ibu Wantiyah yang telah mengizinkan untuk menyelesaikan laporan ini serta memberikan dukungan lahir batin dan doa agar tugas akhir ini dapat berjalan dengan lancar.

Penulis menyadari bahwa dalam penyusunan laporan ini masih terdapat banyak kekurangan, untuk itu penulis mengharap kritik dan saran dari pembaca untuk sempurnanya laporan ini. Semoga dengan ditulisnya laporan ini dapat menjadi sumber ilmu bagi setiap pembaca.

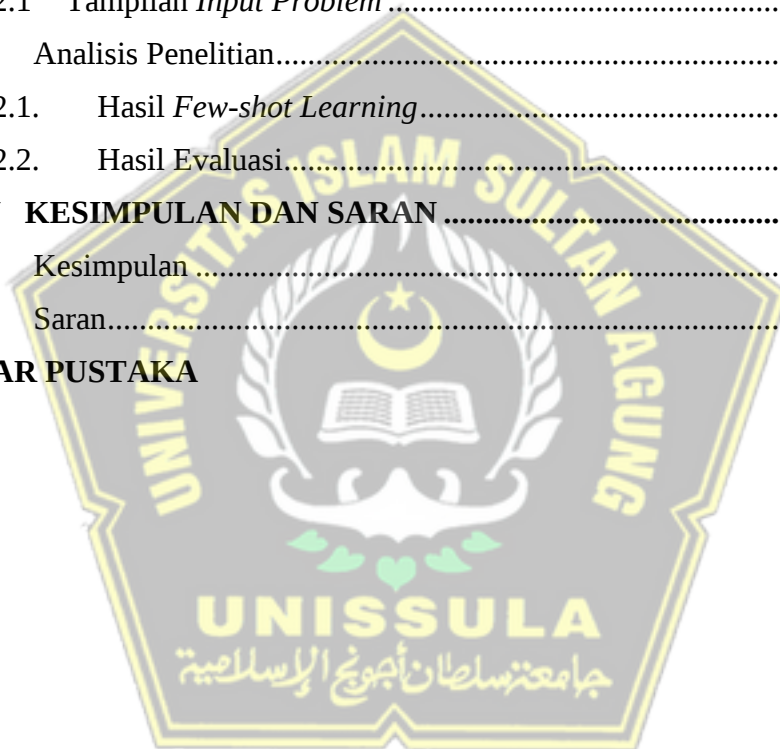
Semarang,.....

Eka Nurul Azizah

DAFTAR ISI

COVER	i
HALAMAN JUDUL	i
LEMBAR PENGESAHAN.....	iii
SURAT PERNYATAAN KEASLIAN TUGAS AKHIR	iv
PERNYATAAN PERSETUJUAN PUBLIKASI KARYA ILMIAH.....	v
KATA PENGANTAR.....	vi
DAFTAR ISI.....	vii
DAFTAR GAMBAR	ix
DAFTAR TABEL.....	x
ABSTRAK.....	xi
BAB I PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Perumusan Masalah	3
1.3 Pembatasan Masalah	3
1.4 Tujuan	3
1.5 Manfaat	3
1.6 Sistematika Penulisan	4
BAB II TINJAUAN PUSTAKA DAN DASAR TEORI	5
2.1 Tinjauan Pustaka	5
2.2 Dasar Teori.....	9
2.2.1 <i>Natural Language Processing (NLP)</i>	9
2.2.2 <i>Large Language Models (LLM)</i>	9
2.2.3 <i>Few-shot Learning</i>	10
2.2.4 Ollama	11
BAB III METODE PENELITIAN	13
3.1 Metode Penelitian.....	13
3.1.1 Studi Literatur	13
3.1.2 Pengumpulan Dataset.....	14
3.1.3 Data Pre Processing.....	16

3.1.4	Pemilihan dan Konfigurasi Model	17
3.1.5	<i>Few-shot Learning</i> Model.....	17
3.1.6	Evaluasi dan Pengujian Model.....	18
3.2	Analisis Sistem.....	19
3.3	Analisis Kebutuhan	21
BAB IV HASIL DAN ANALISIS PENELITIAN		24
4.1	Hasil Penelitian	24
4.1.1	Tampilan Awal Aplikasi	24
4.2.1	Tampilan <i>Input Problem</i>	25
4.2	Analisis Penelitian.....	27
4.2.1.	Hasil <i>Few-shot Learning</i>	28
4.2.2.	Hasil Evaluasi.....	29
BAB V KESIMPULAN DAN SARAN		33
5.1	Kesimpulan	33
5.2	Saran.....	34
DAFTAR PUSTAKA		



DAFTAR GAMBAR

Gambar 2. 1 arsitektur ollama	12
Gambar 3. 1 Alur Penelitian.....	13
Gambar 3. 2 Alur kerja sistem	20
Gambar 4. 1 Tampilan awal UI.....	24
Gambar 4. 2 Tampilan saat Input Problem	25
Gambar 4. 3 Tampilan Output Prediksi	25
Gambar 4. 4 Tampilan Output Prediksi	26
Gambar 4. 5 Tampilan output prediksi.....	26



DAFTAR TABEL

Table 3. 1 Contoh kalimat dari dataset.....	16
Table 4. 1 Evaluasi kalimat Problem-Solution.....	31
Table 4. 2 Hasil Evaluasi Rouge	32



ABSTRAK

Pesatnya pertumbuhan jumlah artikel ilmiah menghadirkan tantangan baru dalam mengekstraksi informasi yang relevan, khususnya dalam mengidentifikasi hubungan antara permasalahan dan solusinya. Penelitian ini bertujuan untuk mengembangkan model prediksi solusi menggunakan Large Language Model (LLM) dengan pendekatan *Few-shot Learning Learning Learning*. Model yang diterapkan adalah Llama 3.2, yang telah disesuaikan dengan dataset hasil ekstraksi dari 100 artikel ilmiah, diklasifikasikan ke dalam empat kategori utama: Problem-Solution, Tantangan-Jawaban, Peluang-Jawaban, dan Kelemahan-Peningkatan. Proses pengolahan data mencakup tahapan pre-processing, seperti case folding, tokenizing, filtering, dan stemming, guna meningkatkan kualitas data sebelum model dilatih. Evaluasi model dilakukan menggunakan metrik ROUGE untuk menilai akurasi prediksi solusi yang dihasilkan. Hasil penelitian menunjukkan bahwa pendekatan *Few-shot Learning Learning Learning* mampu mengenali pola hubungan masalah-solusi dengan lebih efektif dibandingkan metode konvensional. Selain itu, sistem berbasis website juga dikembangkan untuk mempermudah akses dan pemanfaatan model oleh mahasiswa dalam menyelesaikan tugas akhir. Walaupun model menunjukkan kinerja yang baik, tantangan dalam menangani pertanyaan yang berbeda jauh dari contoh yang diberikan masih menjadi kendala yang perlu disempurnakan dalam penelitian mendatang. Kata Kunci: *Few-shot Learning, Large Language Model, Llama 3.2, Prediksi Solusi, ROUGE, Artikel Ilmiah*

ABSTRACT

The rapid growth of scientific articles presents new challenges in extracting relevant information, particularly in identifying the relationship between problems and solutions. This study aims to develop a solution prediction model based on a Large Language Model (LLM) using the *Few-shot Learning Learning Learning* approach. The chosen model, Llama 3.2, is trained with a dataset extracted from 100 scientific articles categorized into four main groups: Problem-Solution, Challenge-Answer, Opportunity-Answer, and Weakness-Improvement. Data processing includes pre-processing steps such as case folding, tokenizing, filtering, and stemming to enhance data quality before model training. Model evaluation is conducted using the ROUGE metric to measure the accuracy of the predicted solutions. The results indicate that the *Few-shot Learning Learning Learning* approach enables the model to recognize problem-solution patterns more efficiently than traditional methods. Additionally, a web-based system was developed to facilitate access and usage for students working on their final projects. Although the evaluation results demonstrate good performance, the model's limitations in handling questions significantly different from the provided examples remain a challenge for future research.

Keywords: *Few-shot Learning, Large Language Model, Llama 3.2, Solution Prediction, ROUGE, Scientific Articles*

BAB I

PENDAHULUAN

1.1 Latar Belakang

Artikel ilmiah merupakan sumber informasi yang sangat penting dalam perkembangan ilmu pengetahuan dan teknologi. Setiap tahun, jumlah publikasi artikel ilmiah terus meningkat secara signifikan. Menurut penelitian yang dilakukan oleh (Bornmann dan Mutz 2015), pertumbuhan publikasi ilmiah global mencapai 4% per tahun, dengan lebih dari 4 juta artikel dipublikasikan pada tahun 2021. Peningkatan yang pesat ini menciptakan tantangan baru dalam hal pengelolaan dan ekstraksi informasi yang relevan dari artikel-artikel tersebut.

Dalam sebuah artikel ilmiah, identifikasi masalah dan solusi merupakan komponen kunci yang sering dicari oleh peneliti dan praktisi (Gao, dkk 2021). dalam penelitiannya mengungkapkan bahwa rata-rata peneliti menghabiskan 12 jam per minggu hanya untuk membaca dan menganalisis artikel ilmiah untuk menemukan solusi yang relevan dengan penelitian mereka. Proses manual ini tidak hanya memakan waktu tetapi juga rentan terhadap kesalahan manusia dan keterbatasan kognitif dalam memproses informasi dalam jumlah besar.

Perkembangan teknologi kecerdasan buatan, khususnya Large Language Models (LLM), telah membuka peluang baru dalam pengolahan dan analisis teks. Penelitian yang dilakukan oleh (Brown, dkk. 2020) menunjukkan bahwa LLM memiliki kemampuan yang menjanjikan dalam memahami konteks dan struktur artikel ilmiah dengan akurasi mencapai 89%. Namun, tantangan utama dalam penggunaan LLM konvensional adalah kebutuhan akan data training yang sangat besar dan resource komputasi yang intensif.

Few-shot Learning muncul sebagai solusi potensial untuk mengatasi keterbatasan tersebut (Samuel dkk. 2022). dalam publikasinya mendemonstrasikan bahwa pendekatan *Few-shot Learning* dapat menghasilkan performa yang kompetitif dengan hanya menggunakan 5-10 contoh training, dibandingkan dengan metode tradisional yang membutuhkan ribuan contoh. Keunggulan ini sangat relevan dalam

konteks analisis artikel ilmiah, dimana setiap bidang memiliki karakteristik dan terminologi yang unik.

Keunggulan utama *Few-shot Learning* dibandingkan fine-tuning terletak pada efisiensi dan fleksibilitasnya. Fine-tuning membutuhkan waktu training yang lama dan komputasi intensif, sementara *Few-shot Learning Learning* dapat beradaptasi secara real-time dengan contoh minimal (Y. Chen dkk. 2022) Dalam konteks prediksi solusi dari masalah artikel ilmiah, *Few-shot Learning Learning* memungkinkan model untuk memahami pola hubungan masalah-solusi dengan lebih efisien dan dapat beradaptasi dengan berbagai domain keilmuan.

Studi komparatif menunjukkan bahwa *Few-shot Learning* menghasilkan performa yang lebih stabil dibandingkan zero-shot learning, terutama dalam tugas yang membutuhkan pemahaman kontekstual yang dalam. Penelitian oleh (P. Liu dkk. 2023) mendemonstrasikan bahwa *Few-shot Learning Learning* mencapai F1-score 15-20% lebih tinggi dibandingkan zero-shot learning dalam tugas analisis teks ilmiah.

Implementasi *Few-shot Learning* dalam konteks prediksi kalimat solusi dari artikel ilmiah menghadapi beberapa tantangan spesifik. Penelitian (P. Liu dkk. 2023) mengidentifikasi bahwa variasi gaya penulisan dan struktur artikel antar disiplin ilmu dapat mempengaruhi akurasi prediksi. Selain itu, kompleksitas bahasa ilmiah dan penggunaan terminologi teknis menambah tingkat kesulitan dalam proses ekstraksi informasi.

Studi komparatif menunjukkan bahwa *Few-shot Learning* menghasilkan performa yang lebih stabil dibandingkan *zero-shot learning*, terutama dalam tugas yang membutuhkan pemahaman kontekstual yang dalam. Integrasi LLM dengan *Few-shot Learning* membuka perspektif baru dalam automasi analisis artikel ilmiah (Keicher 2023). melaporkan peningkatan efisiensi sebesar 75% dalam proses identifikasi solusi ketika menggunakan pendekatan terintegrasi ini. Kemampuan sistem untuk belajar dari contoh terbatas sambil mempertahankan pemahaman kontekstual yang kuat dari LLM menciptakan sinergi yang menjanjikan.

Berdasarkan tantangan dan peluang tersebut, penelitian ini mengusulkan implementasi *Few-shot Learning* dengan model llama 3.2 untuk prediksi kalimat

solusi dari masalah pada artikel ilmiah. Pendekatan ini diharapkan dapat mengatasi keterbatasan dataset berlabel sambil mempertahankan akurasi yang tinggi dalam ekstraksi solusi. Mendemonstrasikan keberhasilan penerapan *Few-shot Learning* dengan model bahasa multilingual, yang dapat menjadi dasar untuk pengembangan sistem yang lebih komprehensif berkontribusi secara signifikan terhadap produktivitas dan kualitas penelitian akademik secara keseluruhan (Lin dkk. 2022)

Implementasi sistem ini diharapkan dapat memberikan kontribusi signifikan dalam meningkatkan efisiensi penelitian ilmiah, mengoptimalkan pemanfaatan pengetahuan yang ada, dan mempercepat inovasi lintas domain. (Y. Liu dkk. 2019) telah meletakkan fondasi yang kuat untuk pengembangan arsitektur transformer dan optimasi model llama 3.2, yang menjadi basis untuk penelitian ini.

1.2 Perumusan Masalah

Bagaimana penerapan pada model berbasis *Large Language Models* (LLM) dapat dioptimalkan untuk memprediksi kalimat solusi dalam artikel secara akurat dan efisien?

1.3 Pembatasan Masalah

1. Model *Large Language Models* (LLM) yang digunakan adalah llama 3.2
2. Penelitian ini hanya berfokus pada kalimat dari artikel ilmiah berbahasa Indonesia
3. Dataset yang digunakan hanya diambil dari artikel ilmiah yang sudah dipublikasikan dan di *review* secara manual oleh peneliti dan divalidasi oleh dosen pembimbing
4. Dataset yang digunakan hanya terdiri dari 4 kategori yaitu, Masalah-Solusi, Tantangan-Jawaban, Peluang-Jawaban, dan Kelemahan-Peringatan.

1.4 Tujuan

Membuat sistem untuk menghasilkan kalimat solusi dari masalah pada artikel ilmiah dengan mengimplementasikan *Few-shot Learning* menggunakan LLM

1.5 Manfaat

1. Membantu pembaca, peneliti, dan akademisi untuk mengakses solusi dalam artikel ilmiah secara lebih cepat dan efisien .

2. Membantu peneliti dalam mengidentifikasi teks solusi dari masalah pada sebuah kalimat secara efisien dan otomatis, sehingga dapat menghemat waktu dalam proses literatur review.

1.6 Sistematika Penulisan

BAB I : PENDAHULUAN

Pada bab 1, penulis mengutarakan urgensi dari penelitian yang diangkat, mulai dari penulisan latar belakang, membuat rumusan masalah, membatasi permasalahan yang dibahas, serta tujuan dan manfaat yang diperoleh, dan sistematika penulisan

BAB II : TINJAUAN PUSTAKA DAN DASAR TEORI

Pada Bab 2, penulis memuat dasar teori yang digunakan, serta rujukan dari penelitian terdahulu yang akan digunakan dalam perancangan sistem, dan membantu penulis untuk memahami teori yang berhubungan dengan *Few-shot Learning*, LLM, dan Ollama selama proses penelitian.

BAB III : METODE PENELITIAN

Pada bab 3, penulis mengungkapkan proses dan tahapan penelitian yang dimulai dari mendapatkan dataset hingga proses pemodelan topik data yang ada.

BAB IV : HASIL DAN ANALISIS PENELITIAN

Pada bab 4, penulis mengungkapkan hasil penelitian, yaitu hasil implementasi *Few-shot Learning* untuk prediksi kalimat Solusi dari masalah pada artikel ilmiah menggunakan LLM

BAB V : KESIMPULAN DAN SARAN

Pada bab 5, penulis memaparkan Kesimpulan dari hasil proses penelitian mulai dari awal sampe akhir

BAB II

TINJAUAN PUSTAKA DAN DASAR TEORI

2.1 Tinjauan Pustaka

Penelitian yang berjudul “*Detecting Hate Speech with GPT-3*” oleh Chiu pada tahun 2022 mengeksplorasi kemampuan model bahasa GPT-3 dalam mendeteksi ujaran kebencian melalui pendekatan *zero-shot*, *one-shot*, dan *Few-shot Learning*. Penelitian ini menggunakan dataset ETHOS (*haTe speech detectiOn dataset*), yang berisi teks daring yang dikategorikan berdasarkan muatan ujaran kebencian. Model GPT-3, yang memiliki kemampuan pembelajaran lintas konteks yang kuat, diuji untuk mengklasifikasikan teks tanpa pelatihan tambahan (*zero-shot*), dengan satu contoh pelatihan (*one-shot*), dan dengan beberapa contoh pelatihan (*Few-shot Learning*). Hasil penelitian menunjukkan bahwa performa model bervariasi tergantung pada pendekatan pembelajaran yang digunakan. Pada *zero-shot learning*, akurasi rata-rata yang dicapai adalah 56%, dengan F1-score sebesar 70%. Sementara itu, untuk *one-shot learning*, performa menurun dengan akurasi rata-rata sebesar 55% dan F1-score sebesar 55%. Sebaliknya, pada *Few-shot Learning*, performa meningkat dengan akurasi rata-rata sebesar 67% dan F1-score sebesar 62%. Hasil ini menyoroti bahwa GPT-3 dapat mengidentifikasi ujaran kebencian dengan efektivitas yang bervariasi berdasarkan jumlah informasi pelatihan yang diberikan. Penelitian ini menegaskan potensi GPT-3 sebagai alat deteksi ujaran kebencian, meskipun terdapat tantangan dalam memastikan hasil yang konsisten dan akurat di semua skenario pembelajaran (Chiu, Collins, dan Alexander 2021).

Penelitian berjudul “*Hate Speech Detection using OpenAI and GPT-3*” yang dilakukan oleh Priya pada tahun 2022 memanfaatkan model GPT-3 untuk mendeteksi ujaran kebencian dan mengklasifikasikan teks menjadi kalimat positif atau negatif. Pendekatan yang digunakan meliputi *zero-shot learning*, *one-shot learning*, *Few-shot Learning*, serta metode *Few-shot Learning approach* dengan kategori campuran (*mixed category*). Dalam *zero-shot learning*, model memprediksi tanpa contoh pelatihan, sementara *one-shot* menggunakan satu

contoh, dan *Few-shot Learning* menggunakan beberapa contoh untuk melatih model memahami konteks lebih baik. Hasil penelitian menunjukkan performa model yang berbeda pada setiap pendekatan. Skor akurasi yang diperoleh adalah 45% untuk *zero-shot learning*, meningkat menjadi 72% pada *one-shot learning*, dan mencapai 80% pada *Few-shot Learning*. Hasil ini mengindikasikan bahwa menambahkan jumlah contoh pelatihan secara signifikan meningkatkan kemampuan GPT-3 dalam memahami dan memprediksi teks berbasis ujaran kebencian, menunjukkan potensi penggunaan model ini dalam deteksi ujaran kebencian dengan pendekatan yang fleksibel (Priya dan Gupta 2022).

Artikel berjudul "*Language Models are Few-shot Learners*" oleh Tom B. Brown dan tim, yang diterbitkan dalam *Advances in Neural Information Processing Systems* (2020), membahas pengembangan dan keunggulan model GPT-3, termasuk kompleksitas, parameter, dan kemampuannya dalam memahami serta menghasilkan bahasa alami. Artikel ini menjelaskan bagaimana GPT-3, dengan 175 miliar parameter, menjadi salah satu model bahasa terbesar dan paling canggih pada masanya. GPT-3 dirancang untuk menangani berbagai tugas tanpa pelatihan khusus, cukup dengan beberapa contoh (*Few-shot Learning*), atau bahkan tanpa contoh (*zero-shot learning*), berkat arsitektur transformer yang sangat efisien. Penulis juga menguraikan kecanggihan teknologi deep learning di balik GPT-3, yang memungkinkan model untuk memahami dan menghasilkan bahasa dengan tingkat kompleksitas tinggi. Meskipun menggunakan istilah teknis yang rumit, artikel ini menggambarkan bagaimana GPT-3 mampu menjawab pertanyaan, menyelesaikan masalah, dan menghasilkan teks dengan kualitas yang mendekati manusia. Keberhasilan tim dalam menciptakan model yang begitu kompleks dan efektif menyoroti kemajuan luar biasa dalam bidang machine learning, khususnya dalam pengembangan model berbasis bahasa alami (Brown, dkk. 2020).

Penelitian berjudul "*LLM-Empowered Chatbots for Psychiatrist and Patient Simulation: Application and Evaluation*" oleh S. Chen dkk. pada tahun 2023 mengeksplorasi penggunaan model bahasa besar (*Large Language Models, LLM*), khususnya ChatGPT, untuk menciptakan chatbot yang mensimulasikan interaksi antara psikiater dan pasien. Penelitian ini memanfaatkan teknik *prompt engineering*

untuk mengarahkan kemampuan ChatGPT agar memberikan respons yang relevan dan kontekstual sesuai dengan kebutuhan simulasi psikiatri. Melalui evaluasi menyeluruh, chatbot dinilai berdasarkan kemampuannya untuk memahami konteks percakapan, memberikan respons yang empatik, dan mendukung skenario klinis. Hasil penelitian menunjukkan bahwa pendekatan berbasis LLM ini memiliki potensi besar dalam aplikasi pelatihan psikiater dan mendukung pasien, meskipun masih ada tantangan seperti memastikan keakuratan respons klinis dan menghindari bias dalam percakapan. Penelitian ini menyoroti bagaimana integrasi teknologi canggih seperti ChatGPT dapat memberikan dampak positif di bidang kesehatan mental (S. Chen dkk. 2023).

Penelitian Jurnal dengan judul “*Image classification based on Few-shot Learning algorithms: a review*” ini membahas tentang klasifikasi gambar menggunakan *Few-shot Learning* (FSL), sebuah pendekatan yang memungkinkan pengklasifikasian gambar meskipun hanya memiliki sedikit data latih. Dalam FSL, terdapat tiga metode utama yang digunakan, yaitu data augmentation, transfer learning, dan meta-learning. Model-model FSL dievaluasi menggunakan beberapa dataset populer, seperti MiniImageNet, Omniglot, CUB-200, TieredImageNet, dan CIFAR-100, dengan metrik penilaian berbasis 5-way 1-shot dan 5-way 5-shot. Berdasarkan hasil eksperimen, model EGNN dan TriNet menunjukkan performa terbaik pada dataset *MiniImageNet*. Namun, FSL masih menghadapi tantangan seperti overfitting, bias dalam dataset, serta keterbatasan dalam penerapannya secara luas. Untuk mengatasi hambatan tersebut, diperlukan pengembangan teknik yang lebih baik agar model memiliki kemampuan generalisasi yang lebih tinggi, interpretabilitas deep learning yang lebih jelas, serta dataset yang lebih representatif. Secara keseluruhan, EGNN dan TriNet menjadi model yang paling unggul dalam klasifikasi *Few-shot Learning* dengan tingkat akurasi tertinggi (Qi, dkk 2024).

Penelitian dengan judul “*Few-shot Learning with Multilingual Generative Language Models*” ini meneliti penggunaan model bahasa generatif multibahasa dalam *Few-shot Learning* untuk berbagai tugas pemrosesan bahasa alami (NLP). Hasil penelitian menunjukkan bahwa model XGLM 7.5B,

yang telah dilatih menggunakan 30 bahasa berbeda, mampu mencapai kinerja optimal dalam *Few-shot Learning* pada lebih dari 20 bahasa, melampaui GPT-3 dengan ukuran yang sebanding dalam tugas seperti penalaran multibahasa, inferensi bahasa alami, dan terjemahan mesin. Model ini mencatat peningkatan akurasi sebesar 7.4% dalam skenario *zero-shot* dan 9.4% dalam skenario *four-shot* untuk penalaran multibahasa, serta meningkatkan inferensi bahasa alami sebesar 5.4% di kedua pengaturan. Pada benchmark FLORES-101, model ini melampaui GPT-3 dalam 171 dari 182 arah terjemahan dengan 32 contoh pelatihan serta melampaui baseline terjemahan berbasis supervisi dalam 45 arah. Studi ini juga membahas berbagai strategi dalam prompting multibahasa, seperti pemanfaatan template dan contoh demonstrasi lintas bahasa untuk meningkatkan performa *Few-shot Learning*. Meskipun demikian, model ini masih mengalami kendala dalam tugas NLP berbahasa Inggris, di mana model yang berfokus pada bahasa Inggris seperti GPT-3 tetap lebih unggul. Secara keseluruhan, penelitian ini menegaskan bahwa model bahasa generatif multibahasa memiliki potensi besar dalam meningkatkan kapabilitas *Few-shot Learning* di berbagai bahasa, terutama pada bahasa dengan sumber daya terbatas (Lin dkk. 2022).

Penelitian dengan judul ” *Zero-Shot Prompting and Few-shot Learning \ Fine-Tuning: Revisiting Document Image Classification Using Large Language Models*” ini mengevaluasi kembali klasifikasi dokumen berbasis citra menggunakan model bahasa besar (LLMs) melalui pendekatan *zero-shot prompting* dan *Few-shot Learning fine-tuning* pada dataset RVL-CDIP. Hasil penelitian menunjukkan bahwa model GPT-4-Vision mencapai akurasi tertinggi sebesar 69,9% dalam skenario *zero-shot* tanpa pelatihan tambahan. Sementara itu, model Mistral-7B dengan *fine-tuning* generatif mampu memperoleh akurasi 83,4% hanya dengan 100 sampel per kelas. Selain itu, metode embedding berbasis teks menggunakan OpenAI-large juga memberikan hasil yang kompetitif dengan akurasi 67,8% pada 1.600 sampel pelatihan. Temuan ini menegaskan bahwa klasifikasi dokumen dapat dilakukan secara efektif meski menggunakan data pelatihan dalam jumlah terbatas, terutama dengan memanfaatkan kombinasi informasi visual dan teks. Ke depannya, pengembangan model multimodal berbasis

LLMs diharapkan dapat lebih meningkatkan kinerja dalam skenario *Few-shot Learning* learning.

2.2 Dasar Teori

2.2.1 *Natural Language Processing (NLP)*

Natural Language Processing (NLP) merupakan kemampuan suatu komputer untuk memahami dan mengolah bahasa manusia, baik dalam bentuk lisan maupun tulisan, sebagaimana digunakan dalam komunikasi sehari-hari. NLP secara sederhana bertujuan untuk *Natural Language Processing (NLP)* adalah area integral dari ilmu komputer antara pembelajaran mesin dan linguistik yang mana komputasi digunakan secara luas (Patel & Prajapati, 2018), yang ditujukan untuk membuat komputer mengerti pernyataan yang ditulis menggunakan bahasa manusia. Pengolahan bahasa alami timbul untuk meringankan pekerjaan user dan untuk memenuhi keinginan terhubung dengan komputer dalam bahasa murni. Karena semua pengguna mungkin tidak fasih dalam bahasa khusus mesin, NLP melayani pengguna yang tidak memiliki cukup waktu untuk mempelajari bahasa baru atau mendapatkan kesempurnaan di dalamnya *Natural Language Processing* berdasarkan basisnya bisa dikelompokkan dalam dua komponen ialah *Natural Language Generation (NLG)* dan *Natural Language Understanding (NLU)* yang mengembangkan tugas guna memahami dan menghasilkan teks. NLP bertindak sebagai pilar dasar untuk pengenalan bahasa, yang digunakan oleh Siri (Apple) dan Google. Hal ini memungkinkan teknologi untuk mengenali teks bahasa alami manusia dan perintah berbasis ucapan dan mencakup dua komponen utama *Natural Language Generation (NLG)* dan *Natural Language Understanding (NLU)*. NLP lebih menyulitkan daripada NLG, karena bahasa alami memiliki struktur dan bentuk yang sangat kompleks, memetakan inputan yang diberikan pengguna dan menganalisis beberapa fitur Bahasa (Dale dkk. 2010).

2.2.2 *Large Language Models (LLM)*

Large Language Model (LLM) adalah salah satu jenis model Bahasa kecerdasan buatan yang termasuk ke dalam *Natural Language Processing (NLP)*. LLM dapat membuat teks dengan kemampuan seperti layaknya manusia. LLM ditandai dengan banyak parameter dan dilatih pada kumpulan teks yang besar, serta

memiliki kemampuan untuk menghasilkan output yang relevan kontekstual dan gramatikal. Sedangkan langchain merupakan kerangka kerja yang digunakan dalam pengembangan aplikasi dengan LLM. Langchain memungkinkan pengembang untuk membangun aplikasi menggunakan LLM untuk meningkatkan penyesuaian, akurat, dan relevansi dari model. Penelitian yang pernah dilakukan oleh Oughuzan Topsakal dan T.Cetin Akinci (2023) tentang *Creating Large Language Model Applications Utilizing Langchain: A Primer on Developing LLM Apps Fast*. Penelitian ini memiliki peran besar dalam dunia pengembangan aplikasi dengan *Large Language Model* (LLM). Penelitian ini juga menggunakan Langchain sebagai kerangka kerja yang memungkinkan pengembang untuk memanfaatkan LLM dengan baik, efektif dan efisien. Langchain juga dapat meningkatkan kualitas dan relevansi output yang dihasilkan. Melalui penelitian ini Topsakal dan Akinci menunjukkan bahwa LLM tidak hanya sebuah alat yang kuat untuk memproses teks, namun juga dapat diterapkan secara luas dalam berbagai konteks, termasuk penulisan kode dan tugas-tugas lainnya dalam bidang kecerdasan buatan (Rizki dkk. 2024).

2.2.3 *Few-shot Learning*

Indikator evaluasi dalam *Few-shot Learning* sangat penting untuk menilai kinerja model yang dilatih dengan data beranotasi terbatas, yang ditugaskan untuk mengklasifikasikan kategori baru secara akurat hanya dengan beberapa sampel berlabel. Biasanya, evaluasi ini dilakukan dalam pengaturan N-way K-shot, di mana setiap tugas terdiri dari support set dan verification set yang mencakup N kategori (way) dan K sampel dalam support set untuk setiap kategori. Akurasi algoritma ditentukan dengan mengujinya pada verification set, mengulangi evaluasi beberapa kali, dan menghitung rata-rata akurasi. Secara lebih spesifik, akurasi klasifikasi pada metode *Few-shot Learning* dihitung berdasarkan proporsi instance yang diklasifikasikan dengan benar dibandingkan dengan ukuran total dataset.

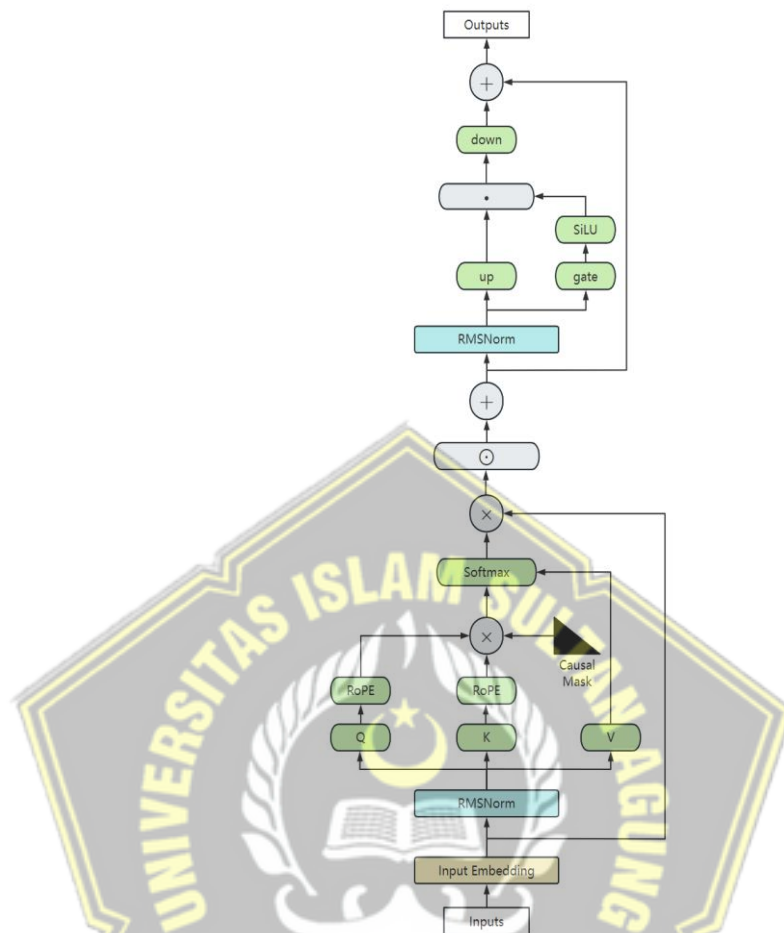
Dalam konteks *Few-shot Learning*, indikator evaluasi penting lainnya adalah efisiensi meta-learning. Efisiensi meta-learning menilai kemampuan model untuk menggeneralisasi pengetahuan di berbagai tugas dan beradaptasi dengan cepat terhadap tugas baru dengan data pelatihan yang minimal. Hal ini sangat

relevan dalam skenario di mana model harus dengan cepat mempelajari konsep atau kategori baru dengan contoh berlabel yang terbatas. Algoritma meta-learning yang efisien dapat memanfaatkan pengetahuan sebelumnya secara efektif dan beradaptasi dengan tugas baru, sehingga meningkatkan kinerja pada tugas klasifikasi *Few-shot Learning*.

Selain akurasi klasifikasi dan efisiensi meta-learning, penting juga untuk mempertimbangkan metrik kinerja lain seperti precision, recall, dan F1-score. Metrik ini memberikan pemahaman yang lebih mendalam tentang kinerja model dengan mempertimbangkan false positive, false negative, serta keseimbangan antara precision dan recall. Dengan mengevaluasi model menggunakan kombinasi akurasi, efisiensi meta-learning, dan metrik relevan lainnya, peneliti dapat menilai efektivitas model secara komprehensif dalam skenario *Few-shot Learning*. (Qi, dkk 2024)

2.2.4 Ollama

Ollama merupakan inovasi terbaru dalam dunia teknologi kecerdasan buatan (AI) yang memungkinkan pengguna menjalankan model bahasa besar (Large Language Models atau LLM) secara lokal di perangkat mereka. Teknologi ini dirancang untuk memenuhi kebutuhan akan akses LLM yang lebih aman, cepat, dan efisien tanpa harus bergantung pada koneksi internet. Ollama menyediakan antarmuka berbasis command-line yang sederhana dan mudah digunakan, sehingga pengguna dapat dengan cepat mengunduh, mengelola, dan mengoperasikan berbagai model AI (Wan dkk. 2024).



Gambar 2. 1 arsitektur ollama

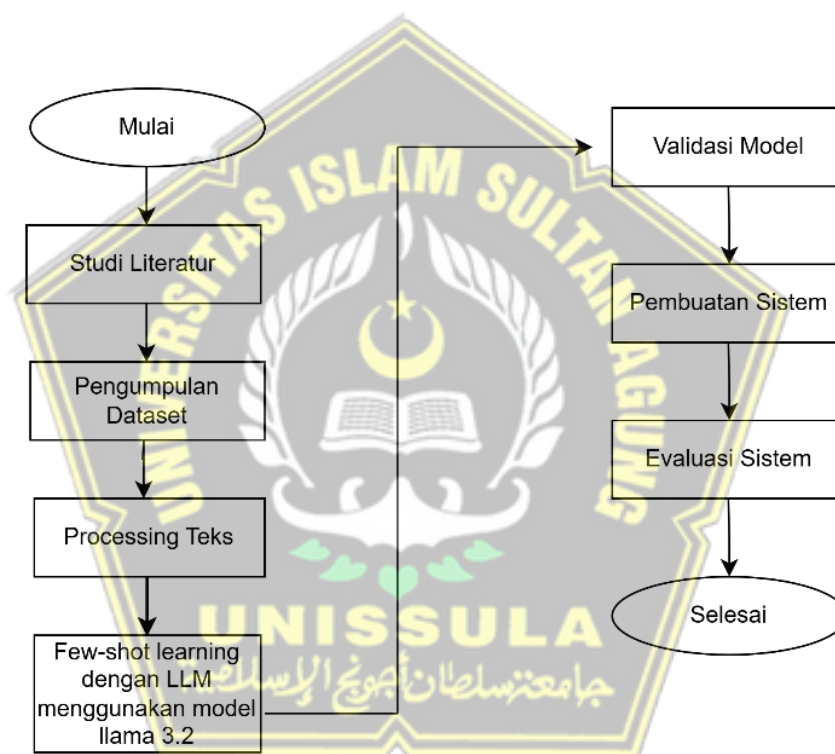
Salah satu model AI yang saya gunakan melalui Ollama adalah Llama 3.2, versi terbaru yang menawarkan peningkatan signifikan dalam performa, kecepatan pemrosesan, dan efisiensi memori. Llama 3.2 dirancang untuk memberikan hasil yang lebih akurat dan responsif, menjadikannya pilihan unggul untuk implementasi lokal AI, termasuk pemrosesan bahasa alami, pengembangan kode, dan aplikasi berbasis AI lainnya. Dengan Ollama dan Llama 3.2, pengguna kini memiliki solusi AI canggih yang dapat diakses kapan saja tanpa perlu koneksi internet

BAB III

METODE PENELITIAN

3.1 Metode Penelitian

Penelitian ini menggunakan pendekatan *Few-shot Learning* menggunakan LLM. Dalam penelitian ini, LLM akan di *training* untuk memprediksi solusi dari kalimat masalah yang terdapat dalam artikel ilmiah. Langkah-langkah yang harus dilakukan dalam penelitian ini adalah sebagai berikut :



Gambar 3. 1 Alur Penelitian

3.1.1 Studi Literatur

Peneliti akan meninjau berbagai sumber literatur, termasuk *e-book*, artikel, jurnal, penelitian sebelumnya seperti tesis dan skripsi, serta informasi dari berbagai situs web yang relevan dengan topik penelitian. Studi literatur ini berfokus pada pemahaman teori mengenai *Few-shot Learning*, (LLM), metode prediksi berbasis teks, serta penerapan model *machine learning* untuk menganalisis solusi, tantangan, peluang, dan kelemahan yang diidentifikasi dalam artikel ilmiah.

3.1.2 Pengumpulan Dataset

Pengumpulan dataset dalam penelitian ini dilakukan secara manual karena belum tersedia dataset yang sesuai dengan kebutuhan penelitian. Dataset ini disusun dengan mereview 100 *artikel*. Artikel yang digunakan diambil dari Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK).

Tahapan pengumpulan data dilakukan sebagai berikut :

a. Identifikasi Permasalahan dalam Artikel

Proses pengumpulan dimulai dengan mengkaji permasalahan yang terdapat di setiap artikel. Setiap kalimat yang memuat permasalahan dan solusinya dicatat untuk membentuk dataset.

b. Kategori Data

Setiap kategori berisi minimal 20 kalimat permasalahan, dan data yang terkumpul sebanyak 99 kalimat. Dataset dikelompokkan ke dalam empat kategori utama, yaitu :

- a) Masalah – Solusi
- b) Tantangan – Jawaban
- c) Peluang – Jawaban
- d) Kelemahan – Peningkatan

c. Pengambilan Kalimat Permasalahan dan Solusi

- a) Jika dalam artikel terdapat lebih dari satu kalimat yang menyatakan permasalahan, semua kalimat tersebut dicatat dan dimasukkan ke dalam dataset.
- b) Jika artikel memuat kalimat permasalahan beserta solusi nya, keduanya dicatat dan dimasukkan ke dalam dataset.
- c) Apabila ditemukan kalimat permasalahan tanpa solusi dalam artikel, solusi dibuat berdasarkan asumsi yang relevan dan diarahkan untuk menyelesaikan masalah tersebut. Asumsi ini kemudian divalidasi oleh dosen pembimbing.

d. Validasi Dataset

Setelah data terkumpul, seluruh dataset diverifikasi oleh dosen pembimbing. Proses validasi bertujuan memastikan bahwa dataset sesuai dengan kebutuhan penelitian dan memiliki kualitas yang memadai.

Sebagai pelengkap, berikut ini disertakan contoh kalimat permasalahan berdasarakan masing-masing kategori.

Tabel 3. 1 Contoh kalimat dari dataset

No	Kategori	<i>Problem</i>	<i>Solution</i>
1	<i>Problem-Solution</i>	Pemesanan barang di Apotek Afdhal masih mengandalkan pembukuan manual, yang sering mengakibatkan kesalahan pendataan barang dan keterlambatan pemesanan	Diusulkan penerapan sistem informasi terintegrasi untuk mengotomatisasi proses pemesanan, sehingga dapat meningkatkan efisiensi dan mengurangi kesalahan dalam pendataan barang
2	Tantangan-Jawaban	Transisi ke rekam medis elektronik menjadi langkah penting, namun juga menghadirkan tantangan besar dalam hal pengelolaan data	Dengan menerapkan metode Agile, rumah sakit memiliki kesempatan untuk mengatasi masalah operasional, meningkatkan ketepatan data, dan menyederhanakan proses pelayanan pasien
3	Peluang-Jawaban	E-commerce yang berkembang pesat membuka peluang bagi fintech untuk memudahkan pembayaran digital di masa mendatang	Dengan meningkatnya penggunaan layanan pembayaran digital, penyedia fintech dapat memanfaatkan tren ini untuk menarik lebih banyak pengguna dengan menyediakan layanan yang aman dan efisien

No	Kategori	<i>Problem</i>	<i>Solution</i>
4	Kelemahan-Peningkatan	Proses pengelolaan data sensus harian rawat inap yang masih manual rentan terhadap ketidakakuratan data	Apabila metode manual terus digunakan, risiko kehilangan data dan ketidakonsistenan informasi akan semakin besar, sehingga penting untuk beralih ke sistem otomatis yang lebih aman dan efisien

3.1.3 Data Pre Processing

Dataset yang telah terkumpul akan melalui tahap pre-processing data untuk meningkatkan kualitas data. Pada tahap ini, data mentah akan diproses melalui serangkaian langkah pembersihan dan standardisasi agar siap digunakan dalam analisis selanjutnya. Berikut adalah tahapan pre-processing data yang akan dilakukan:

1. Case folding atau normalisasi huruf merupakan salah satu tahapan dalam text preprocessing yang bertujuan untuk menyamakan format huruf dengan mengubah seluruh teks menjadi huruf kecil atau huruf besar. Proses ini dilakukan agar teks lebih mudah dibandingkan dan dianalisis tanpa terpengaruh oleh perbedaan penggunaan huruf besar. Case folding yang dilakukan yaitu `text.lower` untuk mengubah teks menjadi huruf kecil semua. Mengubah semua karakter menjadi huruf kecil membantu meningkatkan konsistensi data, membuat analisis dan pemrosesan lebih lanjut menjadi lebih mudah (Salam dkk. 2023).
2. Tokenizing, adalah proses membagi atau memotong teks menjadi kata-kata yang membentuknya. Tujuan dari tokenizing adalah untuk mempermudah proses selanjutnya seperti perhitungan kata, pembobotan kata, dan transformasi data menjadi vektor dengan dimensi tinggi. Dengan tokenizing, data teks dapat diolah dengan lebih mudah dan akurat sebelum dilakukan analisis lebih lanjut (Ariansyah dan Indahyanti 2024).

3. *Filtering* merupakan proses dalam *text preprocessing* setelah *tokenisasi*, *filtering* dilakukan untuk mengambil kata penting hasil tokenisasi. Proses *filtering* dalam membuang kata-kata yang tidak digunakan atau stopword terdapat dalam *bag of words* (Yuniar, dkk 2022).
4. *Stemming* merupakan proses mengubah kata menjadi bentuk dasarnya. *Stemming* dilakukan untuk menyeragamkan bentuk kata. Tujuan dari proses *stemming* adalah menghilangkan imbuhan-imbuhan baik itu berupa *prefiks*, *sufiks*, maupun *konfiks* yang ada pada setiap kata *Stemming* dalam penelitian ini dilakukan berdasarkan aturan morfologi bahasa Indonesia (Yuniar, dkk 2022).

3.1.4 Pemilihan dan Konfigurasi Model

Setelah data tersedia, langkah berikutnya adalah memilih dan mengatur model Llama 3.2 untuk proses *Few-shot Learning*. Model ini dipilih karena kemampuannya dalam memahami serta menghasilkan teks sesuai dengan instruksi yang diberikan. Pada tahap konfigurasi tetap diperhatikan agar model dapat beradaptasi dengan baik tanpa mengalami overfitting.

Selain itu, format input-output juga disesuaikan dengan merancang prompt tertentu agar model lebih memahami tugasnya dalam memberikan solusi berdasarkan pertanyaan yang diberikan. Format prompt yang digunakan adalah:

Dengan pendekatan ini, model memahami bahwa input berupa pertanyaan atau permasalahan yang memerlukan jawaban dalam bentuk solusi, dengan batasan tidak lebih dari dua kalimat.

3.1.5 Few-shot Learning Model

Few-shot Learning merupakan proses penyesuaian model yang telah melalui tahap pelatihan sebelumnya (*pre-trained model*) agar lebih efektif dalam menyelesaikan tugas tertentu atau menangani data tertentu. Model pra-terlatih umumnya dilatih dengan kumpulan data yang luas dan bervariasi, sehingga mampu memahami berbagai jenis data secara umum.

Pada tahap pelatihan, model Llama 3.2 dilatih menggunakan dataset yang telah diproses sebelumnya. 99 data akan di *training* untuk memastikan model mampu mengenali pola hubungan antara kalimat masalah dan solusi dengan baik.

Few-shot Learning model LLama 3.2 pada penelitian ini dilakukan untuk menyesuaikan model agar lebih optimal dalam memprediksi solusi berdasarkan kalimat permasalahan pada artikel ilmiah. Proses ini menggunakan pendekatan pembelajaran *sequence to sequence*, dimana model dilatih untuk memahami hubungan antara input berupa kalimat masalah dan output berupa solusi yang sesuai dengan konteks.

Pelatihan dilakukan dengan strategi evaluasi secara berkala untuk memastikan model dapat beradaptasi dengan baik terhadap data yang digunakan. Selain itu, penyesuaian kecepatan pembelajaran dilakukan untuk menjaga keseimbangan antara stabilitas pelatihan dan kecepatan konvergensi model.

3.1.6 Evaluasi dan Pengujian Model

Setelah model selesai dilatih, dilakukan validasi menggunakan metrik evaluasi seperti ROUGE (*Recall-Oriented Understudy for Gisting Evaluation*). Metrik ini mengukur kesamaan antara solusi yang diprediksi oleh model dengan solusi yang ada dalam dataset. Validasi bertujuan untuk mengukur relevansi, akurasi, dan kualitas prediksi model.

ROUGE (*Recall-Oriented Understudy for Gisting Evaluation*) merupakan metode evaluasi yang banyak digunakan dalam *Natural Language Processing*, khususnya untuk menilai kualitas ringkasan otomatis dan terjemahan mesin. Evaluasi ini dilakukan dengan membandingkan hasil prediksi solusi yang dihasilkan oleh sistem dengan data acuan, yang disebut *ground truth summarization*. ROUGE bekerja menggunakan pendekatan *intrinsic*, yakni menganalisis tingkat kesamaan antara hasil prediksi dan referensi yang diharapkan.

Rumus dasar ROUGE dapat ditulis sebagai berikut :

$$\text{ROUGE} = \frac{\sum_{i=1}^n x_{i \in y}}{y} \times 100\%$$

Keterangan :

- x : jumlah kata relevan yang berhasil diprediksi oleh sistem
- y : total kata dalam referensi atau solusi yang diharapkan.

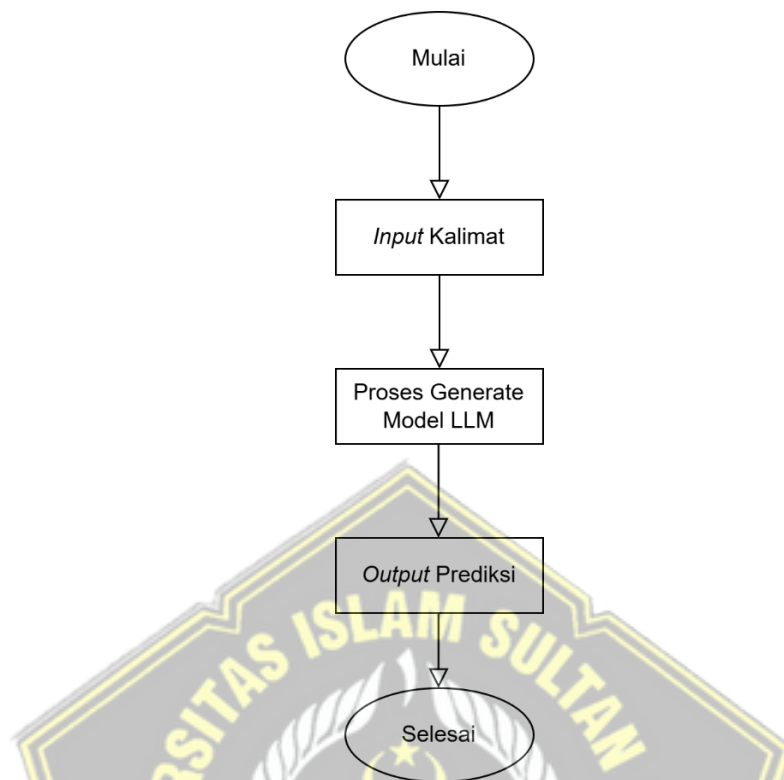
ROUGE menyediakan beberapa jenis metrik yang digunakan sebagai standar untuk menilai keakuratan prediksi suatu sistem. Berikut ini adalah metrik utama yang sering digunakan:

1. ROUGE-1 (unigram): Mengukur tingkat kemiripan satu kata antara prediksi dan referensi. Metrik ini membantu menilai apakah kata-kata penting dalam solusi yang dihasilkan sistem cocok dengan solusi referensi.
2. ROUGE-2: Berfokus pada pasangan kata (bigram) untuk mengevaluasi kesamaan konteks antar kata. ROUGE-2 berguna untuk menilai sejauh mana pasangan kata dalam prediksi mampu mencerminkan konteks atau makna dalam referensi.
3. ROUGE-L (Urutan Umum Terpanjang): Metrik ini menilai kesamaan berdasarkan urutan kata atau struktur kalimat. ROUGE-L digunakan untuk mengukur seberapa cocok struktur atau pola kalimat dalam solusi prediksi dengan solusi referensi, sehingga dapat menilai kealamian dan strukturnya.

Dengan rangkaian metrik ini, ROUGE menjadi alat evaluasi yang ampuh untuk mengukur sejauh mana hasil prediksi suatu sistem mendekati acuan yang diharapkan.

3.2 Analisis Sistem

Pada penelitian ini, penulis akan membuat sistem prediksi solusi berbasis *website* yang bertujuan untuk memberikan layanan informasi tugas akhir kepada mahasiswa di program studi Teknik Informatika UNISSULA. Untuk alurnya akan digambarkan pada gambar 3.2 dengan *flowchart*



Gambar 3. 2 Alur kerja sistem

Berikut adalah tahapan alur kerja sistem :

a. Mulai

Pertama *user* memasuki halaman tampilan awal yang menggambarkan proses dari *website* aplikasi sistem prediksi solusi

b. *Input Kalimat*

Kemudian *user* memasukkan kalimat permasalahan yang terdapat dalam artikel, setelah itu tekan tombol “Enter” atau klik *get solution* setelah *menginputkan text*.

c. Permasalahan diproses

Selanjutnya sistem akan mulai memproses permasalahan yang diberikan oleh *user* dengan mencocokkan data jawaban yang sudah dilatih.

d. *Output Prediksi*

Setelah permasalahan diproses, model akan memberikan *output* prediksi, yaitu kalimat solusi yang relevan dengan masalah yang diberikan. Solusi ini dihasilkan berdasarkan pola-pola yang telah dipelajari oleh model selama proses *Few-shot Learning*.

3.3 Analisis Kebutuhan

Pada tahap analisis kebutuhan, peneliti memeriksa semua perangkat lunak yang diperlukan untuk membuat sistem prediksi solusi ini beroperasi dengan baik dan menghasilkan hasil yang diinginkan. Sistem ini dibuat dengan menggunakan program berikut :

1. *Python*

Python adalah bahasa pemrograman tingkat tinggi yang dirancang agar mudah dibaca dan digunakan dalam berbagai aplikasi, termasuk analisis data dan kecerdasan buatan. Penelitian ini memanfaatkan Python versi 3.12 karena peningkatan kinerjanya, efisiensi penggunaan memori yang lebih baik, serta kompatibilitasnya dengan perpustakaan terbaru yang dibutuhkan dalam pemrosesan bahasa alami (NLP). Untuk menginstalnya, pengguna dapat mengunduhnya dari situs resmi Python <https://www.python.org/downloads/> memilih versi 3.12 sesuai sistem operasi yang digunakan, lalu menjalankan file instalasi sambil mencentang opsi “Add Python to PATH” agar dapat diakses melalui terminal. Setelah proses instalasi selesai, versi Python dapat diverifikasi dengan perintah `python --version`. Dalam penelitian ini, Python digunakan untuk mengeksekusi berbagai skrip yang berkaitan dengan pemrosesan data, pelatihan model, serta pengembangan aplikasi berbasis web.

2. *Library Transformers*

Library Transformers dari Hugging Face menyediakan akses ke berbagai model berbasis transformer yang digunakan dalam tugas NLP, seperti pembuatan teks, klasifikasi, dan ekstraksi informasi. Dalam penelitian ini, Transformers dimanfaatkan untuk memuat model pra-latih serta menyesuaikannya dalam melakukan prediksi solusi dari kalimat permasalahan. Instalasinya dapat dilakukan dengan menjalankan perintah berikut di terminal atau command prompt: `pip install transformers`. Library ini menawarkan antarmuka yang intuitif untuk memuat model besar seperti T5 dan LLaMA, sehingga memungkinkan penelitian ini memanfaatkan teknologi NLP terbaru dalam pemrosesan teks ilmiah.

3. *Library Torch*

Torch berperan sebagai kerangka kerja utama dalam pelatihan dan inferensi model berbasis *Natural Language Processing* (NLP). Dengan Torch, pelatihan model dapat dilakukan secara lebih efisien menggunakan unit pemrosesan grafis (GPU), yang mempercepat proses *Few-shot Learning* model LLaMA 3.2. Instalasi PyTorch dapat dilakukan dengan menjalankan perintah `pip install torch`. Namun, untuk mendukung akselerasi GPU, instalasi harus disesuaikan dengan versi CUDA yang digunakan, sesuai dengan panduan di PyTorch. Setelah berhasil diinstal, PyTorch dapat digunakan dalam kode Python dengan mengimpornya menggunakan `import torch`.

4. *Library Pandas*

Pandas adalah pustaka Python yang digunakan untuk memanipulasi dan menganalisis data dalam format tabel. Dalam penelitian ini, Pandas berfungsi untuk membaca dataset yang berisi pasangan kalimat permasalahan dan solusi dari artikel ilmiah. Instalasi Pandas dapat dilakukan dengan menjalankan perintah `pip install pandas`. Setelah terinstal, pustaka ini dapat digunakan dalam kode Python dengan mengimpornya menggunakan `import pandas as pd`. Pandas bekerja dengan membaca dataset dalam format CSV atau format lainnya, membersihkan data yang tidak relevan, serta melakukan transformasi data agar siap digunakan dalam proses *Few-shot Learning* model NLP.

5. *ROUGE*

Library ROUGE digunakan sebagai metrik evaluasi untuk menilai kualitas prediksi model dengan membandingkannya terhadap solusi referensi. Dalam penelitian ini, ROUGE digunakan untuk mengukur tingkat kesesuaian antara solusi yang dihasilkan oleh model dengan solusi sebenarnya. Instalasi ROUGE dapat dilakukan dengan menjalankan perintah `pip install rouge-score`. Setelah berhasil diinstal, library ini dapat digunakan dalam kode Python dengan mengimpornya menggunakan `from rouge_score import rouge_scorer`. Cara kerja ROUGE melibatkan perhitungan kesamaan n-gram antara teks prediksi dan referensi, menghasilkan skor yang merepresentasikan kualitas prediksi model.

6. *Visual Studio Code (VS Code)*

VS Code adalah editor teks yang populer untuk pengembangan aplikasi karena mendukung berbagai bahasa dan framework pemrograman, serta menyediakan ekstensi yang mempermudah proses pengembangan. Dalam penelitian ini, *VS Code* digunakan untuk menulis dan mengelola kode Python yang digunakan dalam pelatihan model serta pengembangan aplikasi berbasis web. Instalasi dapat dilakukan dengan mengunduhnya melalui situs resmi resmi <https://code.visualstudio.com/>. Setelah terinstal, pengguna dapat menambahkan ekstensi Python untuk mendukung fitur debugging dan menjalankan kode secara langsung di dalam *VS Code*.

7. *Framework Flask*

Flask adalah *framework* Python yang ringan dan bersifat open-source, dirancang untuk membangun aplikasi web dengan fleksibilitas tinggi. Dalam penelitian ini, *Flask* digunakan untuk mengembangkan aplikasi berbasis web yang menampilkan hasil prediksi solusi dari kalimat permasalahan yang diproses oleh model *Few-shot Learning*. *Flask* dipilih karena kemudahannya integrasinya dengan model pembelajaran mesin serta kemampuannya dalam menangani permintaan API untuk pemrosesan data secara real-time. Instalasi *Flask* dapat dilakukan dengan menjalankan perintah `pip install flask`. Setelah berhasil diinstal, *framework* ini dapat digunakan dalam kode Python dengan mengimpornya menggunakan `from flask import Flask`. Dalam penelitian ini, *Flask* bekerja dengan membuat API yang menerima input teks dari pengguna, mengirimkannya ke model NLP, dan mengembalikan hasil prediksi dalam format yang dapat diakses melalui web browser.

BAB IV

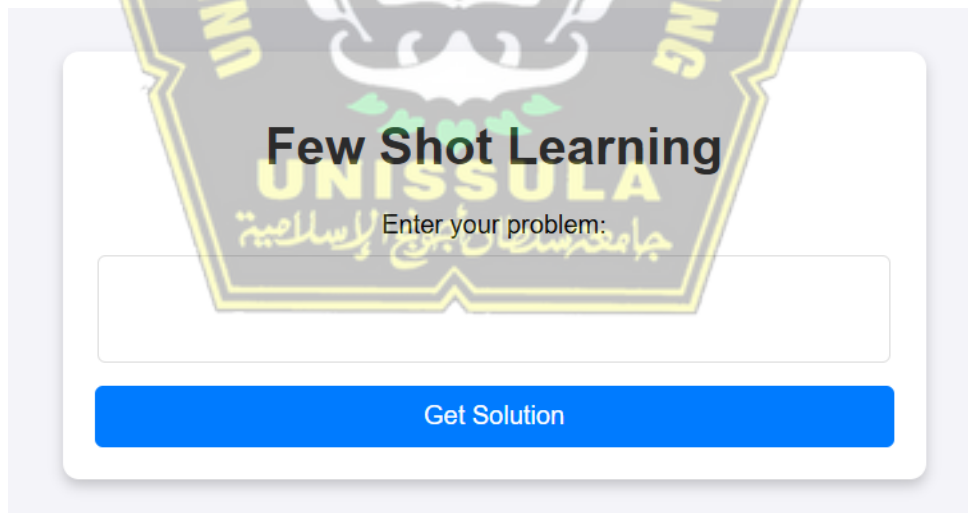
HASIL DAN ANALISIS PENELITIAN

4.1 Hasil Penelitian

4.1.1 Tampilan Awal Aplikasi

Setelah model berhasil diunduh, tahap selanjutnya adalah merancang antarmuka pengguna (UI) guna mempermudah proses memperoleh prediksi solusi dari permasalahan yang dimasukkan. Dalam penelitian ini, penulis menggunakan *Flask*, sebuah framework berbasis *Python* yang ringan dan fleksibel, untuk membangun aplikasi web yang dapat berinteraksi secara langsung dengan model yang telah disesuaikan.

Aplikasi web ini dirancang agar pengguna hanya perlu memasukkan kalimat permasalahan ke dalam kolom input, kemudian sistem akan memprosesnya dan menampilkan hasil prediksi solusi pada halaman *output*. Dengan menerapkan *Flask*, aplikasi ini dapat dijalankan secara lokal maupun diunggah ke server, sehingga memungkinkan akses yang lebih luas bagi pengguna.

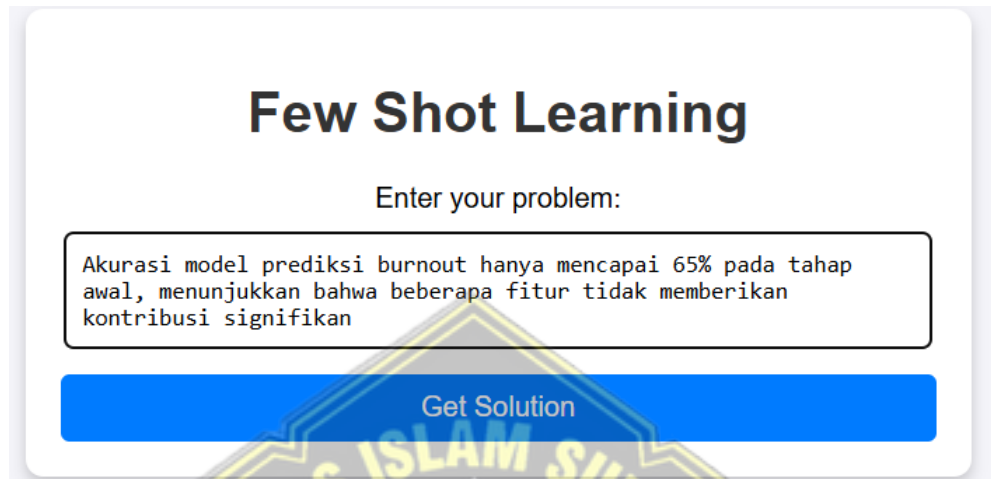


Gambar 4. 1 Tampilan awal UI

Pada Gambar 4.1 Tahap awal *demo* aplikasi *Flask* ini mengintegrasikan model yang telah disesuaikan untuk mempermudah pengguna dalam memperoleh prediksi solusi dalam sistem *Problem-Solution*. Pengguna cukup memasukkan kalimat permasalahan pada kolom *Problem* dan menekan *Get Solution*. Sistem

kemudian akan memproses input tersebut dan menampilkan solusi yang dihasilkan oleh model secara otomatis.

4.2.1 Tampilan *Input Problem*



Few Shot Learning

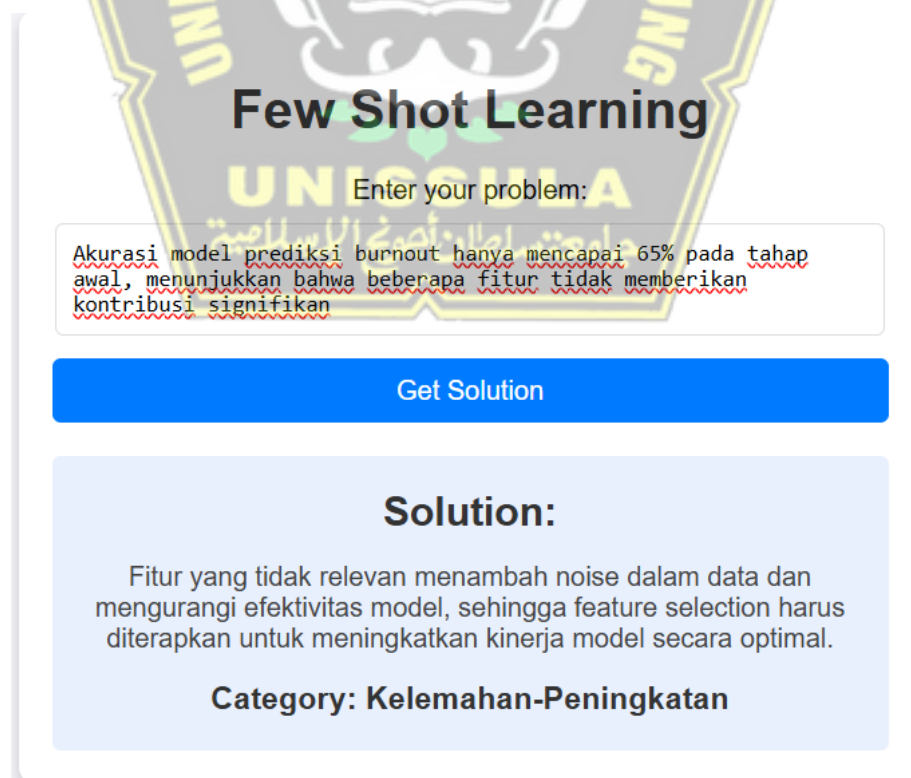
Enter your problem:

Akurasi model prediksi burnout hanya mencapai 65% pada tahap awal, menunjukkan bahwa beberapa fitur tidak memberikan kontribusi signifikan

Get Solution

Gambar 4. 2 Tampilan saat *Input Problem*

Pada Gambar 4.2 tampilan pada saat pengguna memasukkan kalimat masalah untuk mendapatkan prediksi.



Few Shot Learning

Enter your problem:

Akurasi model prediksi burnout hanya mencapai 65% pada tahap awal, menunjukkan bahwa beberapa fitur tidak memberikan kontribusi signifikan

Get Solution

Solution:

Fitur yang tidak relevan menambah noise dalam data dan mengurangi efektivitas model, sehingga feature selection harus diterapkan untuk meningkatkan kinerja model secara optimal.

Category: Kelemahan-Peningkatan

Gambar 4. 3 Tampilan Output Prediksi

Few Shot Learning

Enter your problem:

Penggunaan citra generasi AI dalam konteks hukum dapat merusak kepercayaan pada bukti dan merugikan nilai probatifnya

Get Solution

Solution:

Untuk menjaga integritas bukti hukum, perlu diterapkan regulasi ketat dan metode verifikasi forensik guna membedakan gambar asli dari hasil AI

Category: Problem-Solution

Gambar 4. 4 Tampilan *Output* Prediksi

Few Shot Learning

Enter your problem:

"E-commerce yang berkembang pesat membuka peluang bagi fintech untuk memudahkan pembayaran digital di masa mendatang"

Get Solution

Solution:

Implementasi teknologi blockchain dan internet keamanan dapat membantu meningkatkan keamanan transaksi pembayaran digital dalam e-commerce.

Category: Tantangan-Jawaban

Gambar 4. 5 Tampilan *output* prediksi

Dapat diamati pada gambar 4.3, 4.4 dan 4.5, Sistem menghasilkan prediksi solusi berdasarkan permasalahan yang dimasukkan oleh pengguna. Prediksi yang diberikan dinilai relevan dengan permasalahan yang berkaitan dengan teknologi. Dari empat kali percobaan *demo* aplikasi, sistem terbukti mampu memproses input dan memberikan solusi yang sesuai dengan permasalahan yang diajukan oleh pengguna.

4.2 Analisis Penelitian

Penelitian ini menerapkan *Few-shot Learning*, model LLM untuk menghasilkan prediksi kalimat solusi dari masalah pada artikel ilmiah. Hasil *training* menunjukkan bahwa model terlatih dengan baik, terlihat dari hasil evaluasi menggunakan metrik ROUGE terlihat bahwa relevansi model terhadap acuan jawaban masih rendah, dengan nilai ROUGE-1 sebesar 0,2148, ROUGE-2 sebesar 0,0364, dan ROUGE-L sebesar 0,1479. Nilai tersebut menunjukkan bahwa model kurang mampu menghasilkan teks yang mendekati referensi jawaban secara optimal.

Salah satu tantangan dalam mengembangkan model *Few-shot Learning* LLM untuk menghasilkan prediksi solusi dari permasalahan dalam artikel ilmiah adalah rendahnya relevansi antara keluaran model dan jawaban referensi. Berdasarkan hasil evaluasi menggunakan metrik ROUGE, diketahui bahwa model belum mampu menghasilkan teks yang mendekati jawaban acuan secara optimal, dengan nilai ROUGE-1 sebesar 0,2148, ROUGE-2 sebesar 0,0364, dan ROUGE-L sebesar 0,1479.. Nilai tersebut menunjukkan bahwa model masih mengalami kesulitan dalam mengenali pola bahasa yang sesuai dengan referensi serta memahami keterkaitan antara suatu permasalahan dan solusinya secara efektif.

Meskipun *Few-shot Learning* memungkinkan model untuk belajar dari sejumlah kecil contoh, kinerja model tetap bergantung pada kualitas serta keberagaman data pelatihan yang digunakan. Salah satu kendala utama yang ditemukan adalah model belum mampu melakukan generalisasi dengan baik terhadap contoh terbatas yang diberikan, sehingga sering kali menghasilkan prediksi yang kurang relevan atau tidak sesuai dengan pola solusi yang diharapkan. Beberapa faktor yang dapat memengaruhi performa model antara lain distribusi

data yang tidak seimbang, keterbatasan variasi dalam contoh pelatihan, serta kesulitan model dalam memahami konteks kompleks dengan jumlah contoh yang terbatas.

Selain itu, saat model digunakan dalam proses inferensi, hasil prediksi masih sering kali tidak sesuai dengan input yang diberikan. Salah satu faktor penyebabnya adalah kurangnya adaptasi model terhadap domain tertentu, mengingat *Few-shot Learning* sangat bergantung pada kesesuaian contoh yang diberikan saat inferensi. Jika contoh yang digunakan tidak cukup representatif atau tidak mencakup variasi permasalahan yang luas, maka model cenderung mengalami kesulitan dalam menghasilkan solusi yang akurat. Selain itu, keterbatasan sumber daya komputasi juga dapat berdampak pada performa model, karena model berbasis LLM membutuhkan daya komputasi yang tinggi agar dapat berfungsi secara optimal.

Untuk meningkatkan kinerja model dalam *Few-shot Learning*, beberapa strategi dapat diterapkan, seperti pemilihan contoh yang lebih representatif dan informatif, penerapan teknik prompt engineering guna mengarahkan model dalam menghasilkan output yang lebih akurat. Dengan menerapkan strategi tersebut, diharapkan model dapat menghasilkan prediksi yang lebih relevan dan semakin mendekati jawaban referensi.

4.2.1. Hasil *Few-shot Learning*

Penerapan *Few-shot Learning* menunjukkan bahwa model dapat mengenali pola dari contoh yang diberikan untuk memprediksi jawaban dengan lebih akurat. Pendekatan ini memungkinkan model memahami konteks pertanyaan dan menghasilkan respons yang sesuai meskipun hanya menggunakan sedikit contoh sebagai referensi.

Berdasarkan evaluasi, model mampu memberikan jawaban yang cukup baik, terutama untuk pertanyaan yang memiliki kesamaan dengan contoh yang diberikan. Selain itu, ketika pertanyaan yang diajukan identik dengan data yang tersedia, model dapat memberikan jawaban yang akurat. Namun, efektivitas prediksi tetap dipengaruhi oleh kualitas dan keberagaman contoh yang digunakan.

Jika contoh yang tersedia kurang representatif, model cenderung mengalami kesulitan dalam menghasilkan respons yang tepat.

Secara keseluruhan, *Few-shot Learning* menunjukkan bahwa model mampu memahami serta merespons pertanyaan baru sehingga menghasilkan solusi yang relevan walaupun masih terbilang cukup rendah. Meski demikian, masih terdapat keterbatasan dalam menggeneralisasi pertanyaan yang sangat berbeda dari contoh yang diberikan. Oleh karena itu, pemilihan contoh yang tepat menjadi faktor penting dalam meningkatkan kualitas prediksi model.

4.2.2. Hasil Evaluasi

Untuk memastikan bahwa model yang digunakan mampu memberikan prediksi yang tepat, maka dilakukan evaluasi menggunakan ROUGEscore. Berikut akan ditampilkan tabel perbandingan antara pernyataan prediksi, dan juga tabel hasil evaluasi menggunakan ROUGEscore.

Tabel 4. 1 Evaluasi kalimat *Problem-Solution*

<i>Input</i>	Hasil yang diharapkan	<i>Output</i>	Kesimpulan
Implementasi algoritma Multinomial Naive Bayes membutuhkan data besar untuk hasil yang akurat	Perlu memastikan ketersediaan data yang cukup sebelum mengimplementasikan metode ini	Implementasi algoritma Multinomial Naive Bayes membutuhkan data besar untuk hasil yang akurat karena algoritma ini menggunakan teknik penilaian probabilitas.	Hasil masih Relevan karena hasil yg diharapkan menjelaskan tentang “cukupnya” data sedangkan output nya menjelaskan “besarnya” data
Pedagang keliling saat ini mengalami kesulitan mendapatkan	Mereka dapat memanfaatkan media sosial dan platform online untuk mempromosikan	Pedagang keliling dapat menggunakan media sosial dan aplikasi	Hasil masih Relevan karena hasil yang diharapkan lebih umum

<i>Input</i>	Hasil yang diharapkan	<i>Output</i>	Kesimpulan
pelanggan atau pembeli akibat persaingan dengan penyedia jasa pemesanan makanan secara online	produk mereka serta menawarkan layanan pesan antar yang lebih fleksibel	pemesanan makanan online untuk meningkatkan visibilitas dan menarik pelanggan.	sedangkan outputnya lebih spesifik)
Penerapan sistem otomatis indikator rumah sakit membuka peluang untuk mengurangi kesalahan pencatatan dan mempercepat waktu pengolahan data	Membangun sistem yang mampu memberikan informasi secara real-time, memungkinkan pengambilan keputusan yang lebih cepat dan tepat untuk meningkatkan layanan rumah sakit"	Dengan menggunakan sistem otomatis, data dapat diinput secara langsung dan akurat tanpa adanya kesalahan manual, sehingga mempercepat waktu pengolahan data dan meningkatkan kualitas informasi yang dihasilkan.	Hasil Masih Relevan karena hasil yang diharapkan membahas tentang real-time information system sedangkan outputnya tentang system otomatis
Intrusion Detection System (IDS) ini tidak bekerja maksimal untuk mendeteksi jenis serangan baru	Untuk meningkatkan performa IDS, digunakan pendekatan machine learning, dengan algoritma random forest serta	Implementasikan sistem berbagi pengetahuan dan belajar dari pengalaman melalui algoritma Machine	Hasil Masih Relevan karena keduanya befokus pada peningkatam IDS dengan Machine

<i>Input</i>	Hasil yang diharapkan	<i>Output</i>	Kesimpulan
yang belum pernah terjadi sebelumnya (zero-day)	dataset whitelist dan blacklist	Learning untuk meningkatkan kemampuan IDS dalam mendeteksi serangan zero-day.	Learning tetapi pendekatan yang disebutkan berbeda
Transisi ke rekam medis elektronik menjadi langkah penting, namun juga menghadirkan tantangan besar dalam hal pengelolaan data	Dengan menerapkan metode Agile, rumah sakit memiliki kesempatan untuk mengatasi masalah operasional, meningkatkan ketepatan data, dan menyederhanakan proses pelayanan pasien.	Mengadopsi sistem manajemen data yang terintegrasi dan aman dapat membantu mengatasi tantangan yang dihadapi selama transisi ke rekam medis elektronik.	Hasil Kurang Relevan karena hasil yang diharapkan membahas tentang penerapan metode Agile sedangkan outputnya membahas system manajemen data

Berdasarkan tabel 4.1 di atas, hasil evaluasi menunjukkan bahwa dari lima sampel yang diuji, empat di antaranya berhasil menghasilkan prediksi yang masih relevan dengan input yang diberikan. Namun, satu sampel lainnya menghasilkan prediksi yang kurang sesuai. Pengujian ini dilakukan dengan menginputkan lima kalimat masalah dari total 20 data evaluasi yang digunakan dalam model yang telah dilatih. Oleh karena itu, dapat disimpulkan bahwa tidak semua kalimat masalah yang dimasukkan akan selalu menghasilkan prediksi solusi yang relevan.

Tabel 4. 2 Hasil evaluasi ROUGE

ROUGE-1	ROUGE-2	ROUGE-L
0.2148	0.0364	0.1479

Berdasarkan hasil evaluasi menggunakan ROUGE, terlihat bahwa model *Few-shot Learning* memiliki tingkat kesesuaian yang cukup baik dengan acuan. ROUGE-1 memiliki skor 0.2148, yang menunjukkan bahwa sekitar 21,48% unigram pada hasil prediksi telah sesuai dengan acuan. Sementara itu, ROUGE-2 memiliki skor 0.0364, yang berarti hanya 3,64% bigram yang diprediksi sesuai dengan referensi. Nilai ini lebih rendah dibandingkan ROUGE-1 karena konformitas bigram lebih sulit dicapai. Selanjutnya, ROUGE-L memiliki skor 0.1479, yang menunjukkan bahwa urutan kata dalam prediksi memiliki kesesuaian sebesar 14,79% dengan referensi.

Dari hasil evaluasi ini, dapat disimpulkan bahwa model *Few-shot Learning* mampu menghasilkan prediksi yang cukup relevan. Namun, karena metode ini bekerja dengan jumlah data yang terbatas, masih terdapat potensi perbaikan, terutama dalam memastikan bahwa contoh-contoh yang digunakan dalam *Few-shot Learning* memiliki kualitas yang representatif dan beragam. Selain itu, eksplorasi metode augmentasi data, seperti paraphrasing atau back translation, dapat membantu meningkatkan kualitas input tanpa perlu melakukan *fine-tuning* secara penuh. Dengan strategi ini, model diharapkan dapat menghasilkan prediksi yang lebih akurat meskipun dengan data yang terbatas.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil penelitian dapat disimpulkan bahwa implementasi *Few-shot Learning* menggunakan model LLaMA 3.2 dalam skenario *Few-shot Learning* menunjukkan tingkat kesesuaian yang cukup baik terhadap referensi, meskipun masih terdapat beberapa keterbatasan dalam akurasi prediksi. Skor ROUGE-1 sebesar 0.2148 menunjukkan bahwa model berhasil menangkap yang sesuai dengan referensi. Namun, skor ROUGE-2 yang lebih rendah, yaitu 0.0364, mengindikasikan bahwa model masih mengalami kesulitan dalam mengenali hubungan antar kata pada tingkat bigram. Sementara itu, skor ROUGE-L sebesar 0.1479 menunjukkan bahwa model memiliki kesesuaian urutan kata dengan referensi.

Secara keseluruhan, model LLaMA 3.2 dalam konfigurasi *Few-shot Learning* mampu menghasilkan prediksi yang cukup relevan, meskipun masih terdapat ruang untuk perbaikan, terutama dalam memahami struktur dan keterkaitan antar kata yang lebih kompleks. Untuk meningkatkan akurasi prediksi, perlu dilakukan optimasi dengan memilih contoh yang lebih representatif dan bervariasi dalam *Few-shot Learning*, serta mempertimbangkan metode augmentasi data, seperti paraphrasing atau back translation, guna meningkatkan kualitas input tanpa memerlukan *fine-tuning* secara menyeluruh. Dengan penerapan strategi ini, model diharapkan dapat menghasilkan prediksi yang lebih akurat meskipun dalam kondisi data yang terbatas.

5.2 Saran

Untuk meningkatkan akurasi kinerja model dalam memprediksi solusi kalimat dari masalah pada sebuah artikel ilmiah, ada beberapa saran yang dapat dipertimbangkan untuk pengembangan lebih lanjut:

1. Menggunakan contoh data yang lebih beragam dan sesuai dengan konteks permasalahan agar model dapat mengenali pola dengan lebih baik.
2. Menerapkan teknik augmentasi data, seperti paraphrasing atau back translation, untuk memperluas cakupan data tanpa perlu menambah jumlah data pelatihan secara signifikan.
3. Melakukan eksperimen dengan jumlah contoh berbeda guna menemukan konfigurasi terbaik yang memberikan hasil prediksi paling optimal.
4. Menambahkan data yang lebih relevan dengan bidang yang diteliti untuk membantu model memahami konteks secara lebih mendalam.
5. Menggunakan metrik lain, seperti *BLEU* atau *BERTScore*, untuk memperoleh penilaian yang lebih komprehensif terhadap kualitas prediksi model.

DAFTAR PUSTAKA

- Ariansyah, Achmad, dan Uce Indahyanti. 2024. "Fitur Ekstraksi pada Pemodelan Topik Menggunakan Metode Latent Dirichlet Allocation pada Peristiwa Kebocoran Data." (2): 1–24.
- Bornmann, Lutz, dan Rüdiger Mutz. 2015. "Growth rates of modern science: A bibliometric analysis based on the number of publications and cited references." *Journal of the Association for Information Science and Technology* 66(11): 2215–22.
- Brown, Tom B., Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, dkk. 2020. "Language models are *Few-shot Learning* learners." *Advances in Neural Information Processing Systems* 2020-Decem.
- Brown, Tom B, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Gretchen Krueger, dkk. 2020. "Language Models are *Few-shot Learning* Learners." (Juli 2023): 1877–1901.
- Chen, Siyuan dkk. 2023. "LLM-empowered Chatbots for Psychiatrist and Patient Simulation: Application and Evaluation."
- Chen, Yanda dkk. 2022. "Meta-learning via Language Model In-context Tuning." *Proceedings of the Annual Meeting of the Association for Computational Linguistics* 1: 719–30.
- Chiu, Ke-Li, Annie Collins, dan Rohan Alexander. 2021. "Detecting Hate Speech with GPT-3." (March): 1–29.
- Dale, Robert, Brent L Iverson, Peter B Dervan, dan M R F Hadi. 2010. "Natural Language Processing menggunakan Metode Fuzzy String Matching pada Chatbot berbasis Web pada Aplikasi Whatsapp." *Handbook of Natural Language Processing, Second Edition*: 7823–30.
http://digilib.uinsby.ac.id/id/eprint/51730%0Ahttp://digilib.uinsby.ac.id/51730/2/MuhammadRifqyFakhrulHadi_H06217015.pdf.
- Gao, Tianyu, Adam Fisch, dan Danqi Chen. 2021. "Making pre-trained language models better *Few-shot Learning* learners." *ACL-IJCNLP 2021 -*

59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing, Proceedings of the Conference: 3816–30.

Keicher, Philip. 2023. “Machine Learning in Top Physics in the ATLAS and CMS Collaborations.” : 4–9. <http://arxiv.org/abs/2301.09534>.

Lin, Xi Victoria dkk. 2022. “Few-shot Learning Learning with Multilingual Generative Language Models.” *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing, EMNLP 2022: 9019–52*.

Liu, Pengfei dkk. 2023a. “Pre-train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing.” *ACM Computing Surveys* 55(9).

———. 2023b. “Pre-train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing.” *ACM Computing Surveys* 55(9): 1–46.

Liu, Yinhan dkk. 2019. “RoBERTa: A Robustly Optimized BERT Pretraining Approach.” (1). <http://arxiv.org/abs/1907.11692>.

Priya, dan Sachin Gupta. 2022. “Hate Speech Detection using OpenAI and GPT-3.” *International Journal of Emerging Technology and Advanced Engineering* 12(5): 132–38.

Qi, Qiao, Azlin Ahmad, dan Wang Ke. 2024. “Image classification based on Few-shot Learning Learning learning algorithms: a review.” *Indonesian Journal of Electrical Engineering and Computer Science* 35(2): 933–43.

Rizki, Febrian dkk. 2024. “KLIK: Kajian Ilmiah Informatika dan Komputer Implementasi Question Answering Berbasis Chatbot Telegram Pada Tafsir Al-Jalalain Menggunakan Langchain dan LLM.” *Media Online* 4(5): 2464–72.

Salam, Rizky Rahman dkk. 2023. “Analisis Sentimen Terhadap Bantuan Langsung Tunai (BLT) Bahan Bakar Minyak (BBM) Menggunakan Support Vector Machine.” *MALCOM: Indonesian Journal of Machine Learning and Computer Science* 3(1): 27–35.

Samuel, David dkk. 2022. “Direct parsing to sentiment graphs.” *Proceedings of the*

- Annual Meeting of the Association for Computational Linguistics* 2: 470–78.
- Wan, Yao dkk. 2024. “Deep Learning for Code Intelligence: Survey, Benchmark and Toolkit.” *ACM Computing Surveys* 1(1).
- Yuniar, Eka, Dwi Safiroh Utsalinah, dan Dian Wahyuningsih. 2022. “Implementasi Scrapping Data Untuk Sentiment Analysis Pengguna Dompot Digital dengan Menggunakan Algoritma Machine Learning.” *Jurnal Janitra Informatika dan Sistem Informasi* 2(1): 35–42.

