

Penerapan IndoBERT untuk Pencarian Semantik Tafsir Al-Qur'an Berdasarkan Tafsir Al-Misbah

Muh Taslim Sultan ¹, Faisal Akib ^{2*}, Muhammad Hasrul Hasanuddin ³

^{1,2,3} Jurusan Teknik Informatika, Universitas Islam Negeri Alauddin Makassar

¹ 60200120060@uin-alauddin.ac.id, ² faisal@uin-alauddin.ac.id, ³ muhammad.hasrul@uin-alauddin.ac.id

Diajukan: 12/07/2026 | Direvisi: 16/08/2026 | Diterima: 17/08/2026 | Diterbitkan: 18/08/2026

Abstract

Conventional Qur'anic exegesis retrieval systems that rely on keyword matching often return less relevant results because they are unable to capture the semantic meaning and context of users' queries. This study aims to develop a semantic search system for Tafsir Al-Misbah using the IndoBERT model within the CRISP-DM framework. The dataset consists of 6,236 Indonesian-language tafsir summaries processed through data cleaning, case folding, normalization, and tokenization. The IndoBERT model was then fine-tuned to generate semantic representations of user queries and document contents, while FAISS was employed to accelerate retrieval using cosine similarity. System performance was evaluated using Mean Average Precision (mAP), Mean Reciprocal Rank (MRR), and Normalized Discounted Cumulative Gain (nDCG). The evaluation results achieved Top-10 scores of 0.2128 for mAP, 0.2681 for MRR, and 0.2792 for nDCG, while the Top-50 scores reached 0.2301, 0.2803, and 0.3481, respectively. These results indicate adequate retrieval performance in the domain of Qur'anic exegesis and demonstrate that the proposed semantic search approach based on IndoBERT and FAISS is capable of providing contextually relevant retrieval results for users' queries.

Keywords: FAISS, IndoBERT, Semantic Search, Tafsir Al-Misbah, website

Abstrak

Pencarian tafsir Al-Qur'an yang masih mengandalkan pencocokan kata kunci sering kali menghasilkan informasi yang kurang relevan karena belum mampu memahami makna dan konteks pertanyaan pengguna. Penelitian ini bertujuan mengembangkan sistem *semantic search* untuk Tafsir Al-Misbah menggunakan model IndoBERT dengan kerangka kerja CRISP-DM. Dataset yang digunakan terdiri atas 6.236 ringkasan tafsir berbahasa Indonesia yang diproses melalui tahapan *cleaning*, *case folding*, normalisasi, dan tokenisasi. Selanjutnya, model IndoBERT *fine-tune* untuk menghasilkan representasi semantik dari *query* dan dokumen, sedangkan proses pencarian dipercepat menggunakan FAISS dengan perhitungan *cosine similarity*. Kinerja sistem dievaluasi menggunakan metrik *Mean Average Precision* (mAP), *Mean Reciprocal Rank* (MRR), dan *Normalized Discounted Cumulative Gain* (nDCG). Hasil pengujian menunjukkan nilai Top-10 sebesar mAP 0,2128, MRR 0,2681, dan nDCG 0,2792, sedangkan pada Top-50 diperoleh mAP 0,2301, MRR 0,2803, dan nDCG 0,3481. Nilai tersebut menunjukkan performa yang memadai dalam menghasilkan pencarian yang relevan pada domain tafsir Al-Qur'an, sehingga pendekatan *semantic search* berbasis IndoBERT dan FAISS mampu memahami konteks pertanyaan pengguna dengan lebih baik dibandingkan pencocokan kata kunci.

Kata kunci: FAISS, IndoBERT, Semantic Search, Tafsir Al-Misbah, website

This work is licensed under a Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International License (CC BY-NC-SA 4.0). Copyright (C) Author's.



1. PENDAHULUAN

Perkembangan teknologi informasi telah mengubah cara masyarakat memperoleh pengetahuan keagamaan, termasuk dalam mengakses tafsir Al-Qur'an. Berbagai kitab tafsir kini tersedia dalam bentuk digital sehingga informasi dapat diperoleh dengan lebih mudah melalui perangkat elektronik. Kondisi tersebut memberikan kemudahan bagi pengguna untuk mencari penjelasan suatu ayat sesuai dengan kebutuhan. Namun, semakin banyaknya koleksi tafsir yang tersedia juga menimbulkan tantangan dalam menemukan informasi yang benar-benar sesuai dengan maksud pertanyaan pengguna, terutama ketika pencarian dilakukan menggunakan bahasa alami. Oleh karena itu, diperlukan sistem pencarian yang tidak hanya mampu mengakses dokumen secara cepat, tetapi juga memahami konteks pertanyaan agar hasil yang diberikan lebih relevan [1].

Meskipun berbagai platform digital telah menyediakan fasilitas pencarian tafsir Al-Qur'an, sebagian besar masih mengandalkan pendekatan berbasis *keyword matching*. Pendekatan ini bekerja dengan mencocokkan kata yang dimasukkan pengguna dengan kata yang terdapat pada dokumen, sehingga hasil pencarian sangat bergantung pada kesamaan leksikal. Akibatnya, sistem sering kali tidak mampu memahami maksud pertanyaan ketika pengguna menggunakan sinonim, parafrasa, atau susunan kalimat

yang berbeda dengan redaksi yang terdapat pada dokumen tafsir. Kondisi tersebut menyebabkan hasil pencarian menjadi kurang relevan dan berpotensi menyulitkan pengguna dalam memperoleh informasi yang sesuai dengan konteks pertanyaannya [2].

Untuk mengatasi keterbatasan tersebut, pendekatan *semantic search* dikembangkan dengan memanfaatkan representasi semantik sehingga sistem mampu memahami hubungan makna antara pertanyaan pengguna dan dokumen, tidak hanya berdasarkan kesamaan kata. Berbeda dengan pencarian berbasis *keyword*, *semantic search* mempertimbangkan konteks, sinonim, serta hubungan antarkonsep dalam teks sehingga dapat menghasilkan dokumen yang lebih relevan terhadap maksud pengguna. Seiring perkembangan *Natural Language Processing* (NLP), pendekatan ini semakin banyak diterapkan pada berbagai sistem temu kembali informasi karena terbukti mampu meningkatkan kualitas hasil pencarian, khususnya pada dokumen yang bersifat tekstual dan kaya konteks [3].

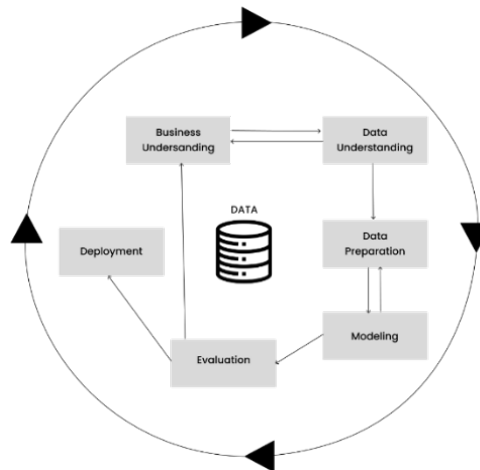
Berbagai penelitian telah menerapkan pendekatan *semantic search* dan model berbasis Transformer untuk meningkatkan kualitas pencarian informasi pada domain keislaman. Penelitian oleh Liza [4] mengembangkan sistem pencarian semantik tafsir Al-Qur'an menggunakan algoritma *cosine similarity* yang dipadukan dengan *Large Language Models* (LLMs), sehingga mampu meningkatkan relevansi hasil pencarian berdasarkan makna pertanyaan pengguna. Sementara itu, Malik [5] mengembangkan sistem *SoulScripture* yang memanfaatkan model BERT sebagai chatbot berbasis Al-Qur'an dan Hadis untuk memberikan bimbingan spiritual melalui pemahaman semantik terhadap pertanyaan pengguna. Meskipun demikian, kedua penelitian tersebut belum secara khusus memanfaatkan model IndoBERT yang dioptimalkan untuk bahasa Indonesia pada proses pencarian tafsir Al-Qur'an, serta belum mengintegrasikan representasi semantik dengan mekanisme pengindeksan vektor menggunakan FAISS untuk mendukung proses temu kembali informasi yang efisien. Kondisi tersebut menunjukkan bahwa penerapan *semantic search* berbasis IndoBERT pada pencarian tafsir Al-Qur'an berbahasa Indonesia masih terbatas dan masih memerlukan pengembangan lebih lanjut.

Berdasarkan kondisi tersebut, penelitian ini memiliki dua kontribusi utama. Pertama, penelitian ini mengembangkan sistem pencarian semantik tafsir Al-Qur'an berbahasa Indonesia dengan memanfaatkan model IndoBERT yang telah di-*fine-tune* sehingga mampu menghasilkan representasi semantik yang lebih sesuai dengan karakteristik bahasa Indonesia. Kedua, penelitian ini mengintegrasikan FAISS sebagai mekanisme pengindeksan vektor untuk mempercepat proses pencarian tanpa mengurangi relevansi hasil. Dengan demikian, sistem yang dikembangkan diharapkan mampu menyajikan ayat, terjemahan, dan ringkasan tafsir Al-Misbah yang relevan berdasarkan konteks pertanyaan pengguna, sekaligus memberikan kontribusi terhadap pengembangan sistem temu kembali informasi pada domain literatur keislaman. [6].

Penelitian ini memiliki kebaruan pada penerapan *semantic search* berbasis IndoBERT yang diintegrasikan dengan FAISS untuk pencarian tafsir Al-Qur'an berbahasa Indonesia. Berbeda dengan pendekatan pencarian berbasis *keyword*, sistem yang dikembangkan memanfaatkan representasi semantik untuk memahami hubungan makna antara pertanyaan pengguna dan dokumen tafsir, serta pengindeksan vektor untuk meningkatkan efisiensi proses pencarian. Tujuan penelitian ini adalah mengembangkan sistem pencarian yang mampu menyajikan ayat, terjemahan, dan tafsir Al-Qur'an yang relevan berdasarkan konteks pertanyaan pengguna. Hasil penelitian ini diharapkan dapat menjadi alternatif dalam pengembangan sistem temu kembali informasi berbasis *semantic search*, khususnya pada domain literatur keislaman.

2. METODE PENELITIAN

Penelitian ini menggunakan pendekatan eksperimental untuk mengembangkan dan mengevaluasi sistem *semantic retrieval* tafsir Al-Qur'an berbasis model IndoBERT. Proses pengembangan model dan implementasi sistem dilakukan menggunakan kerangka *Cross-Industry Standard Process for Data Mining* (CRISP-DM) yang terdiri atas enam tahapan, yaitu *business understanding*, *data understanding*, *data preparation*, *modeling*, *evaluation*, dan *deployment* [7]. Kerangka kerja tersebut digunakan sebagai panduan sistematis dalam pengolahan data, pembentukan model, evaluasi kinerja, hingga implementasi sistem ke dalam website. Melalui tahapan tersebut, sistem diharapkan mampu menghasilkan pencarian tafsir Al-Qur'an yang relevan berdasarkan konteks pertanyaan pengguna.



Gambar 1. CRISP-DM Life Cycle

2.1. Dataset Penelitian

Dataset yang digunakan pada penelitian ini merupakan data sekunder yang telah dikompilasi menjadi korpus penelitian berupa kumpulan ringkasan Tafsir Al-Misbah berbahasa Indonesia. Dataset disusun dalam format Microsoft Excel dan digunakan sebagai basis pengetahuan pada sistem *semantic search*. Setiap entri dataset memuat identitas ayat, teks ayat Al-Qur'an, terjemahan bahasa Indonesia, serta ringkasan Tafsir Al-Misbah sehingga dapat digunakan pada tahapan *preprocessing*, pembentukan *embedding*, dan *semantic retrieval*.

Secara keseluruhan dataset penelitian terdiri atas 6.236 entri, di mana setiap entri merepresentasikan satu ayat Al-Qur'an beserta identitas ayat, teks ayat, terjemahan bahasa Indonesia, dan ringkasan Tafsir Al-Misbah. Sebelum digunakan pada proses pembentukan model, dataset melalui tahapan *preprocessing* yang meliputi *cleaning*, *case folding*, normalisasi, dan tokenisasi untuk menyeragamkan format penulisan serta meningkatkan kualitas data sehingga siap digunakan pada proses pembentukan representasi semantik menggunakan model IndoBERT.

Tabel 1. Karakteristik Dataset

Karakteristik	Keterangan
Sumber Data	Data Sekunder
Jumlah Data	6.236 Baris Ringkasan Tafsir
Representasi Data	1 entri = 1 ayat Al-Qur'an
Bahasa	Indonesia
Bentuk Data	Teks

Contoh dataset yang digunakan sebagai korpus penelitian disajikan pada Tabel 2.

Tabel 2. Contoh Dataset

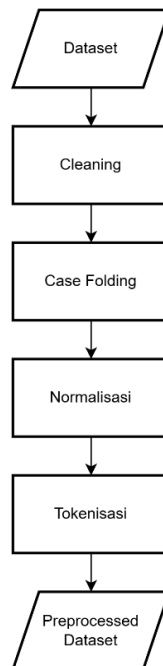
Post_Title	Isi_Ayat	Terjemah_Indo	Tafsir_Almisbah
Tafsir Surah Al-Fatihah Ayat 1 (QS.1:1)	بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ	[1] Dengan menyebut nama Allah Yang Maha Pemurah lagi Maha Penyayang.	[1] Surat ini dimulai dengan menyebut nama Allah--satu-satunya Tuhan yang berhak disembah--Yang memiliki seluruh sifat kesempurnaan dan tersucikan dari segala bentuk kekurangan. Dialah Pemilik rahmah

Post_Title	Isi_Ayat	Terjemah_Indo	Tafsir_Almisbah
			(sifat kasih) yang tak habis-habisnya, Yang menganugerahkan segala macam kenikmatan, baik besar maupun kecil.
Tafsir Surah Al-Fatihah Ayat 2 (QS.1:2)	٢ الْحَمْدُ لِلَّهِ رَبِّ الْعَالَمِينَ	[2] Segala puji bagi Allah, Tuhan semesta alam.	[2] Segala puja dan puji kita persembahkan kepada Allah semata, karena Dialah Yang menciptakan dan memelihara seluruh makhluk.
Tafsir Surah Al-Fatihah Ayat 3 (QS.1:3)	٣ الرَّحْمَنُ الرَّحِيمُ	[3] Maha Pemurah lagi Maha Penyayang.	[3] Dia adalah Pemilik dan Sumber rahmah yang tak pernah putus memberikan segala kenikmatan, baik besar maupun kecil.

Berdasarkan contoh pada Tabel 3, setiap entri dataset memuat informasi mengenai identitas ayat, terjemahan dalam bahasa Indonesia, dan ringkasan tafsir yang digunakan sebagai sumber informasi pada proses *semantic retrieval*. Struktur data tersebut memungkinkan sistem memahami hubungan semantik antara pertanyaan pengguna dan isi tafsir sehingga dapat menghasilkan hasil pencarian yang lebih relevan.

2.2. Preprocessing Dataset

Tahap *preprocessing* merupakan proses awal yang dilakukan untuk mempersiapkan dataset sebelum digunakan pada proses pembentukan representasi semantik menggunakan model IndoBERT. Tahapan ini bertujuan untuk meningkatkan kualitas data dengan menghilangkan ketidakkonsistenan penulisan, menyeragamkan format teks, serta mempersiapkan struktur data agar sesuai dengan kebutuhan pemrosesan bahasa alami (*Natural Language Processing*). Dengan demikian, data yang dihasilkan menjadi lebih bersih, terstruktur, dan mampu meningkatkan efektivitas proses pencarian semantik [8]. Adapun tahap *preprocessing* yang diterapkan dalam penelitian ini ditunjukkan pada gambar 2.



Gambar 2. Tahapan *Preprocessing* Dataset

Berdasarkan Gambar 2, proses *preprocessing* dilakukan secara bertahap mulai dari dataset awal hingga menghasilkan dataset yang siap digunakan pada proses pembentukan *embedding*. Setiap tahapan memiliki fungsi yang berbeda, mulai dari membersihkan data, menyeragamkan penulisan teks, memperbaiki kata yang tidak baku, hingga memecah teks menjadi satuan token. Seluruh proses tersebut bertujuan untuk menghasilkan representasi teks yang lebih konsisten sehingga dapat diproses secara optimal oleh model IndoBERT pada tahap berikutnya.

2.2.1. Dataset

Dataset yang digunakan berupa ringkasan tafsir Al-Qur'an dalam bahasa Indonesia yang berasal dari Tafsir Al-Misbah. Dataset ini menjadi sumber data utama yang diproses sebelum digunakan pada sistem *semantic search*.

2.2.2. Cleaning

Tahap *cleaning* dilakukan untuk menghapus karakter, simbol, dan spasi berlebih yang tidak diperlukan. Proses ini bertujuan menghasilkan data yang lebih bersih dan konsisten.

2.2.3. Case Folding

Case folding dilakukan dengan mengubah seluruh huruf menjadi huruf kecil (*lowercase*). Tahap ini bertujuan menyeragamkan format penulisan agar memudahkan proses pengolahan teks.

2.2.4. Normalisasi

Tahap normalisasi bertujuan mengubah kata yang tidak baku menjadi bentuk baku. Proses ini membantu mengurangi variasi penulisan yang dapat memengaruhi hasil pencarian.

2.2.5. Tokenisasi

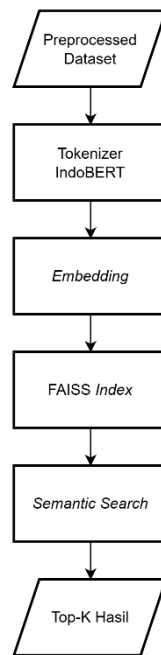
Tokenisasi merupakan proses memecah teks menjadi satuan kata (token). Hasil tokenisasi digunakan sebagai masukan pada proses pembentukan representasi semantik oleh model IndoBERT.

2.2.6. Dataset Hasil Preprocessing

Hasil dari seluruh tahapan *preprocessing* adalah dataset yang telah bersih, terstruktur, dan memiliki format yang konsisten. Dataset tersebut selanjutnya digunakan sebagai masukan pada proses pembentukan *embedding* menggunakan model IndoBERT dalam sistem *semantic search*.

2.3. Arsitektur Semantic Search

Arsitektur *semantic search* dirancang untuk menghasilkan pencarian tafsir Al-Qur'an berdasarkan makna pertanyaan pengguna. Sistem memanfaatkan IndoBERT untuk membentuk representasi semantik dan FAISS untuk mempercepat proses pencarian dokumen.



Gambar 3. Tahapan *Semantic Search*

Berdasarkan Gambar 3, proses *semantic search* dimulai dari dataset hasil *preprocessing*, kemudian dilakukan tokenisasi menggunakan *tokenizer* IndoBERT, pembentukan *embedding*, pengindeksan menggunakan FAISS, hingga menghasilkan dokumen tafsir yang paling relevan (Top-K hasil).

2.3.1. *Preprocessed Dataset*

Dataset hasil *preprocessing* digunakan sebagai masukan utama dalam proses pembentukan representasi semantik.

2.3.2. *Tokenizer IndoBERT*

Teks diubah menjadi token sesuai dengan format masukan model IndoBERT sebelum diproses lebih lanjut.

2.3.3. *Embedding*

Model IndoBERT menghasilkan representasi vektor (*embedding*) yang menggambarkan makna semantik setiap dokumen.

2.3.4. *FAISS Index*

Embedding yang dihasilkan diindeks menggunakan FAISS untuk mempercepat proses pencarian dokumen.

2.3.5. *Semantic Search*

Sistem menghitung kemiripan semantik antara *query* pengguna dan dokumen yang telah diindeks untuk menemukan hasil yang paling relevan.

2.3.6. *Top K-Hasil*

Dokumen dengan nilai kemiripan tertinggi ditampilkan sebagai hasil pencarian beserta ayat, terjemahan, dan tafsir yang sesuai.

2.4. Fine Tuning IndoBERT

Fine-tuning IndoBERT dilakukan untuk menyesuaikan model pralatih (*pre-trained model*) dengan karakteristik dataset tafsir Al-Qur'an yang digunakan pada penelitian ini. Proses ini bertujuan meningkatkan kemampuan model dalam memahami hubungan semantik antara pertanyaan pengguna dan dokumen tafsir sehingga menghasilkan pencarian yang lebih relevan. Konfigurasi parameter *fine-tuning* yang digunakan dalam penelitian ini ditunjukkan pada Tabel 3.

Tabel 3. Parameter *Fine-Tuning*

Parameter	Nilai
Model	Indobert-base-p1
Epoch	5
Batch Size	16
Learning Rate	2e-5
Loss Function	Multiple Negative Ranking Loss

Dataset pasangan *query-document* dibagi menjadi data pelatihan (*training set*), data validasi (*validation set*), dan data pengujian (*test set*) dengan proporsi 80%:10%:10%. *Training set* digunakan untuk proses *fine-tuning* model IndoBERT, sedangkan *validation set* digunakan pada setiap *epoch* untuk memantau kinerja model serta menentukan model terbaik. Setelah proses pelatihan selesai, model dievaluasi menggunakan *test set* yang tidak digunakan selama proses pelatihan sehingga nilai *Mean Average Precision* (mAP), *Mean Reciprocal Rank* (MRR), dan *Normalized Discounted Cumulative Gain* (nDCG) mencerminkan kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya.

2.5. Metode Evaluasi Kinerja *Semantic Search*

Evaluasi kinerja *semantic search* dilakukan untuk mengukur kemampuan sistem dalam mengembalikan dokumen tafsir yang relevan berdasarkan *query* yang diberikan. *Query* evaluasi diperoleh dari dataset yang telah dikompilasi dan dipasangkan dengan dokumen tafsir yang relevan (*ground truth*) pada tahap persiapan data. Selanjutnya, dataset dibagi menjadi data pelatihan, validasi, dan pengujian untuk mendukung proses *fine-tuning* dan evaluasi model. Pada tahap pengujian, setiap *query* digunakan sebagai masukan (*input*) ke sistem, kemudian hasil temu kembali dibandingkan dengan dokumen relevan (*ground truth*) untuk mengukur tingkat relevansi hasil pencarian. Kinerja sistem selanjutnya dievaluasi menggunakan metrik *Mean Average Precision* (mAP), *Mean Reciprocal Rank* (MRR), dan *Normalized Discounted Cumulative Gain* (nDCG), sehingga diperoleh gambaran mengenai kemampuan sistem dalam menghasilkan hasil pencarian yang akurat dan terurut sesuai tingkat relevansinya.

Mean Average Precision (mAP) digunakan untuk mengukur rata-rata ketepatan sistem dalam mengembalikan dokumen yang relevan untuk seluruh *query*.

$$mAP = \frac{1}{Q} \sum_{q=1}^Q AP(q) \quad (1)$$

Mean Reciprocal Rank (MRR) digunakan untuk mengukur posisi dokumen relevan pertama yang ditemukan oleh sistem.

$$MRR = \frac{1}{Q} \sum_{i=1}^Q \frac{1}{rank_i} \quad (2)$$

Normalized Discounted Cumulative Gain (nDCG) digunakan untuk mengevaluasi kualitas peringkat dokumen berdasarkan tingkat relevansinya.

$$nDCG@K = \frac{nDCG@K}{IDCG@k} \quad (3)$$

3. HASIL PENELITIAN DAN PEMBAHASAN

Bagian ini menyajikan hasil penelitian yang diperoleh dari setiap tahapan pengembangan sistem *semantic search* berbasis IndoBERT untuk pencarian tafsir Al-Qur'an. Hasil penelitian disajikan secara

bertahap mulai dari *preprocessing* data, *fine-tuning* model, evaluasi kinerja sistem, hingga implementasi dalam *platform* website. Selanjutnya, setiap hasil dianalisis untuk menjelaskan performa sistem serta keterkaitannya dengan tujuan penelitian.

3.1. Hasil *Preprocessing*

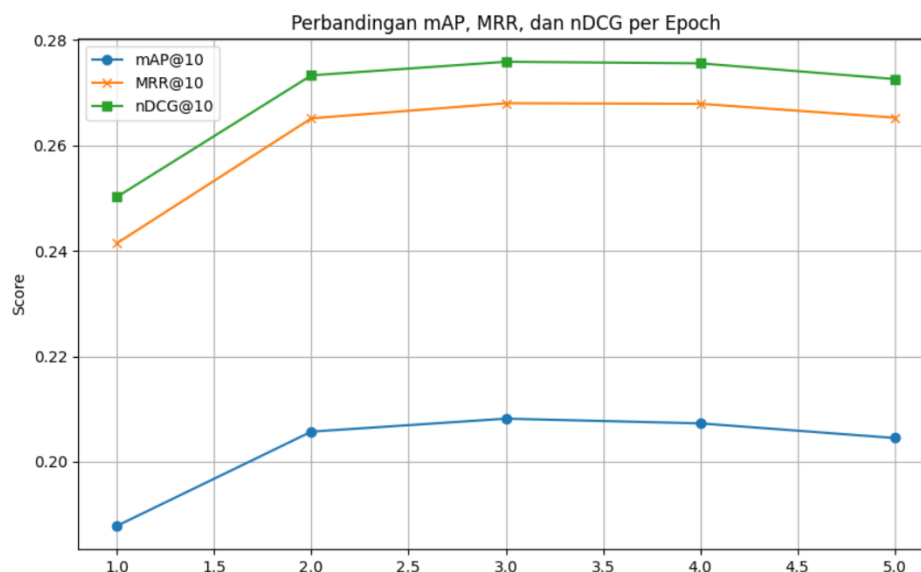
Tahap *preprocessing* menghasilkan dataset yang lebih bersih dan konsisten sehingga siap digunakan pada proses pembentukan *embedding* menggunakan model IndoBERT. Proses ini dilakukan melalui tahapan *cleaning*, *case folding*, normalisasi, dan tokenisasi untuk mengurangi ketidakakonsistenan penulisan serta meningkatkan kualitas data. Contoh hasil *preprocessing* ditampilkan pada Tabel 4.

Tabel 4. Perbandingan Hasil *Preprocessing* data

Tahapan	Sebelum	Sesudah
Cleaning	[2] Segala puja dan puji kita persembahkan kepada Allah semata, karena Dialah Yang menciptakan dan memelihara seluruh makhluk.	Segala puja dan puji kita persembahkan kepada Allah semata karena Dialah Yang menciptakan dan memelihara seluruh makhluk
Case Folding	[2] Segala puja dan puji kita persembahkan kepada Allah semata, karena Dialah Yang menciptakan dan memelihara seluruh makhluk.	segala puja dan puji kita persembahkan kepada allah semata karena dialah yang menciptakan dan memelihara seluruh makhluk
Normalisasi	[2] Segala puja dan puji kita persembahkan kepada Allah semata, karena Dialah Yang menciptakan dan memelihara seluruh makhluk.	segala puja dan puji kita persembahkan kepada allah semata karena dialah yang menciptakan dan memelihara seluruh makhluk
Tokenisasi	[2] Segala puja dan puji kita persembahkan kepada Allah semata, karena Dialah Yang menciptakan dan memelihara seluruh makhluk.	[segala, puja, dan, puji, kita, persembahkan, kepada, allah, semata, karena, dialah, yang, menciptakan, dan, memelihara, seluruh, makhluk]

3.2. Hasil *Fine-Tuning*

Proses *fine-tuning* dilakukan untuk menyesuaikan model IndoBERT dengan dataset tafsir Al-Qur'an sehingga mampu menghasilkan representasi semantik yang lebih baik. Selama proses pelatihan, model mempelajari karakteristik dan hubungan makna antar teks agar dapat memahami konteks pertanyaan



Gambar 4. Evaluasi Hasil *Fine-tuning*

pengguna secara lebih akurat. Kinerja model dievaluasi pada setiap *epoch* menggunakan metrik *Mean Average Precision* (mAP), *Mean Reciprocal Rank* (MRR), dan *Normalized Discounted Cumulative Gain* (nDCG) untuk memantau perkembangan performa selama proses pelatihan. Hasil evaluasi pada setiap epoch ditunjukkan pada Gambar 4.

Berdasarkan Gambar 4, nilai mAP, MRR, dan nDCG mengalami peningkatan dari *epoch* pertama hingga *epoch* ketiga. Setelah itu, performa model cenderung stabil pada *epoch* keempat dan mengalami sedikit penurunan pada *epoch* kelima. Hasil tersebut menunjukkan bahwa model telah mencapai performa terbaik pada *epoch* ketiga dengan nilai mAP sebesar 0,208, MRR sebesar 0,268, dan nDCG sebesar 0,276, sehingga model pada *epoch* tersebut digunakan pada tahap *semantic search*.

3.3. Hasil Evaluasi *Semantic Search*

Evaluasi *semantic search* dilakukan untuk mengukur kemampuan model dalam mengembalikan dokumen tafsir yang relevan berdasarkan *query* pada data pengujian (*test set*). Proses evaluasi dilakukan dengan membandingkan hasil temu kembali sistem terhadap dokumen relevan (*ground truth*) menggunakan metrik *Mean Average Precision* (mAP), *Mean Reciprocal Rank* (MRR), dan *Normalized Discounted Cumulative Gain* (nDCG). Pengukuran dilakukan pada skenario Top-10 dan Top-50 untuk mengetahui kualitas hasil pencarian berdasarkan tingkat relevansi dokumen yang berhasil ditemukan. Hasil evaluasi kinerja sistem disajikan pada Tabel 5.

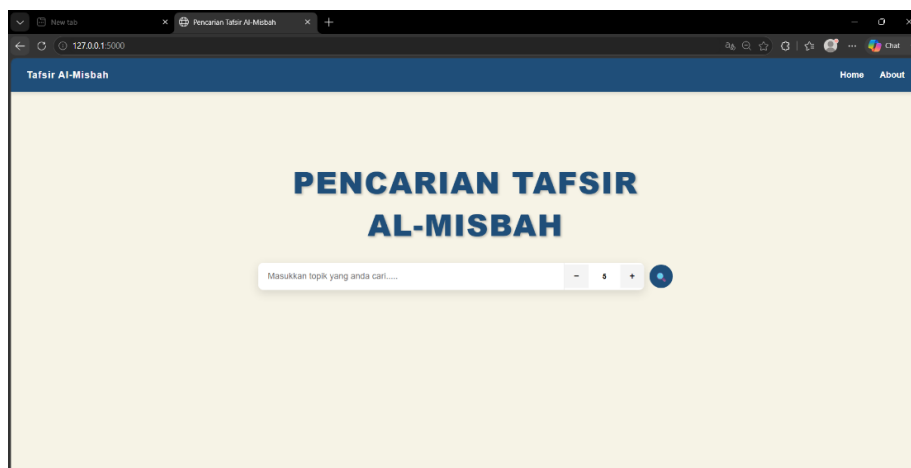
Tabel 5. Hasil Evaluasi Seluruh *Epoch*

Top K	mAP	MRR	nDCG
K = 10	0.2128	0.2681	0.2792
K = 20	0.2224	0.2756	0.3096
K = 50	0.2301	0.2803	0.3481

Berdasarkan Tabel 5, sistem memperoleh nilai mAP sebesar 0,2128, MRR 0,2681, dan nDCG 0,2792 pada skenario Top-10. Sementara itu, pada skenario Top-50 nilai evaluasi meningkat menjadi 0,2301 untuk mAP, 0,2803 untuk MRR, dan 0,3481 untuk nDCG. Peningkatan nilai pada seluruh metrik menunjukkan bahwa penambahan jumlah dokumen yang dievaluasi memberikan peluang lebih besar bagi sistem untuk menemukan dokumen yang relevan. Secara keseluruhan, nilai tersebut menunjukkan performa yang memadai pada domain pencarian semantik tafsir Al-Qur'an, di mana sistem mampu memahami hubungan makna antara *query* dan dokumen secara lebih baik dibandingkan pendekatan berbasis pencocokan kata kunci.

3.4. Hasil Pencarian Tafsir

Hasil implementasi sistem ditampilkan melalui antarmuka website yang dirancang untuk memudahkan pengguna dalam melakukan pencarian tafsir Al-Qur'an menggunakan bahasa alami. Sistem



Gambar 5. Tampilan *Home Screen*

memproses pertanyaan pengguna melalui mekanisme *semantic search* dan menampilkan hasil pencarian berupa ayat, terjemahan, serta tafsir yang memiliki tingkat relevansi tertinggi.

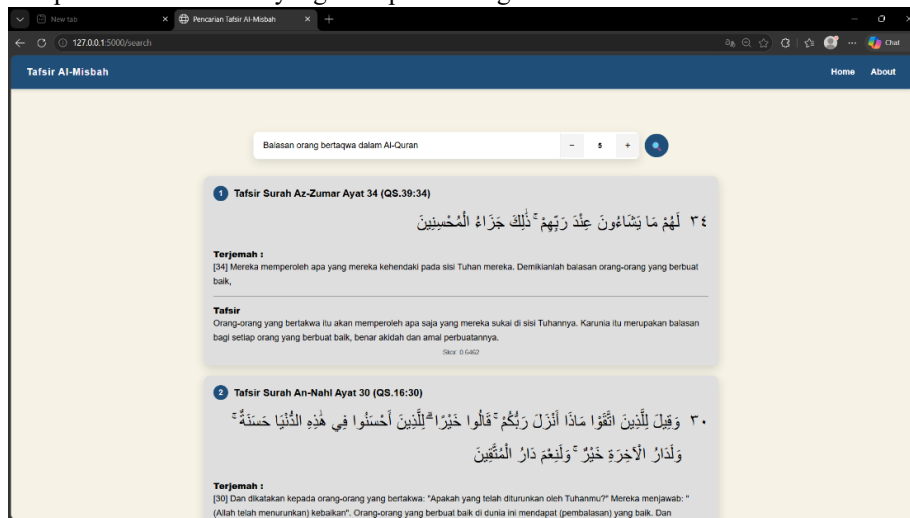
Gambar 5 menunjukkan halaman utama website yang menyediakan kolom input bagi pengguna untuk memasukkan pertanyaan terkait tafsir Al-Qur'an. Antarmuka tersebut dirancang untuk mendukung interaksi berbasis bahasa alami sehingga pengguna dapat mengajukan pertanyaan secara lebih fleksibel. *Query* yang dimasukkan kemudian diproses menggunakan mekanisme *semantic search* untuk menghasilkan hasil pencarian yang relevan sesuai dengan konteks pertanyaan pengguna.

Gambar 6 menampilkan hasil pencarian setelah pengguna mengajukan pertanyaan. Sistem menampilkan ayat Al-Qur'an, terjemahan, dan tafsir yang memiliki tingkat relevansi tertinggi berdasarkan hasil *semantic search*. Hasil pencarian diperoleh melalui perhitungan kemiripan semantik antara *query* pengguna dan representasi *embedding* dokumen yang telah diindeks menggunakan FAISS. Dengan pendekatan ini, sistem mampu menyajikan informasi yang relevan meskipun terdapat perbedaan kosakata antara pertanyaan pengguna dan isi dokumen tafsir.

3.5. Pembahasan

Berdasarkan hasil evaluasi, penerapan model IndoBERT pada sistem *semantic search* mampu meningkatkan kemampuan pencarian dalam menemukan tafsir Al-Qur'an yang relevan dengan pertanyaan pengguna. Hal ini ditunjukkan oleh nilai *Mean Average Precision* (mAP), *Mean Reciprocal Rank* (MRR), dan *Normalized Discounted Cumulative Gain* (nDCG) yang menunjukkan bahwa sistem mampu menempatkan dokumen relevan pada peringkat hasil pencarian secara efektif. Ketiga metrik tersebut merupakan ukuran yang umum digunakan dalam evaluasi sistem *information retrieval* karena mampu mengukur relevansi dokumen berdasarkan ketepatan, posisi hasil pencarian, dan kualitas peringkat dokumen yang dikembalikan [9].

Penggunaan IndoBERT berkontribusi terhadap peningkatan kualitas pencarian karena mampu membentuk representasi semantik yang mempertimbangkan konteks kata dalam suatu kalimat, sehingga



Gambar 6. Tampilan Hasil Pencarian

sistem tidak hanya bergantung pada kesamaan kata kunci [10]. Representasi *embedding* yang dihasilkan memungkinkan pencarian menemukan dokumen tafsir yang memiliki kedekatan makna dengan pertanyaan pengguna meskipun menggunakan kosakata yang berbeda. Karakteristik ini menjadikan pendekatan *semantic search* lebih efektif dibandingkan pencarian berbasis leksikal, khususnya pada tugas *information retrieval* yang memerlukan pemahaman konteks untuk menghasilkan dokumen yang relevan [11].

Selain kemampuan IndoBERT dalam membentuk representasi semantik, penggunaan FAISS (*Facebook AI Similarity Search*) berperan penting dalam meningkatkan efisiensi proses pencarian dokumen. FAISS membangun indeks vektor sehingga proses pencarian dokumen dengan tingkat kemiripan tertinggi dapat dilakukan secara cepat tanpa membandingkan *query* terhadap seluruh koleksi dokumen

secara langsung. Pendekatan ini sangat sesuai untuk sistem *semantic search* karena mampu menangani kumpulan *embedding* berdimensi tinggi dengan tetap mempertahankan kualitas hasil pencarian. Oleh karena itu, integrasi IndoBERT dan FAISS tidak hanya meningkatkan relevansi hasil pencarian, tetapi juga mendukung performa sistem dalam memberikan respons yang lebih cepat pada website [12].

Implementasi sistem pada website menunjukkan bahwa pengguna dapat mengajukan pertanyaan menggunakan bahasa alami tanpa harus menyesuaikan kata kunci tertentu. Sistem kemudian memproses *query* melalui tahapan *semantic search* untuk menghasilkan ayat Al-Qur'an, terjemahan, dan tafsir yang paling relevan sesuai konteks pertanyaan [13]. Pendekatan ini meningkatkan kualitas interaksi antara pengguna dan sistem karena proses pencarian tidak hanya mempertimbangkan kesamaan kata, tetapi juga hubungan semantik antara *query* dan dokumen. Hasil tersebut sejalan dengan penelitian yang menyatakan bahwa penerapan *semantic search* pada website sistem pencarian mampu meningkatkan relevansi respons dan memberikan pengalaman pengguna yang lebih baik dibandingkan pendekatan berbasis pencocokan kata kunci [14].

Dibandingkan dengan pendekatan pencarian berbasis *keyword* atau TF-IDF, sistem yang dikembangkan pada penelitian ini mampu menghasilkan pencarian yang lebih kontekstual karena mempertimbangkan representasi semantik antara *query* dan dokumen [15]. Pendekatan tersebut memungkinkan sistem menemukan tafsir yang relevan meskipun terdapat perbedaan kosakata antara pertanyaan pengguna dan isi dokumen. Temuan ini sejalan dengan penelitian yang menunjukkan bahwa metode *dense retrieval* berbasis *transformer* memiliki kemampuan yang lebih baik dalam memahami hubungan semantik dibandingkan metode *lexical retrieval*, terutama pada tugas *information retrieval* yang memerlukan pemahaman konteks. Dengan demikian, integrasi IndoBERT dan FAISS memberikan kontribusi terhadap peningkatan relevansi hasil pencarian sekaligus mendukung proses temu kembali informasi yang lebih efektif [16].

Temuan ini sejalan dengan penelitian yang menunjukkan bahwa metode *dense retrieval* berbasis *transformer* memiliki kemampuan yang lebih baik dalam memahami hubungan semantik dibandingkan metode *lexical retrieval*, terutama pada tugas *information retrieval* yang memerlukan pemahaman konteks. Dengan demikian, integrasi IndoBERT dan FAISS memberikan kontribusi terhadap peningkatan relevansi hasil pencarian sekaligus mendukung proses temu kembali informasi yang lebih efektif [16]. Sebagai bukti empiris atas keunggulan pendekatan ini, analisis kualitatif dapat dilihat dari hasil pencarian menggunakan *query* "balasan orang bertaqwa dalam Al-Quran" seperti yang ditunjukkan pada Gambar 6. Jika menggunakan pendekatan *keyword matching* konvensional, sistem akan sangat bergantung pada kemunculan kata eksak secara bersamaan, sehingga berpotensi gagal mengembalikan dokumen relevan jika dokumen tersebut menggunakan padanan kata lain seperti "ganjaran", "surga", atau "kebaikan". Sebaliknya, model IndoBERT usulan mampu memahami makna secara kontekstual dan berhasil mengembalikan dokumen yang sangat relevan, seperti Tafsir Surah Az-Zumar Ayat 34 dan An-Nahl Ayat 30. Pada dokumen-dokumen tersebut, sistem berhasil menangkap makna semantik dari kata "karunia", "kebaikan", dan "pembalasan yang baik" yang berkorelasi erat dengan maksud *query* pengguna. Hal ini membuktikan bahwa integrasi IndoBERT mampu mengatasi kendala variasi kosakata yang menjadi kelemahan utama pada pencarian berbasis leksikal.

4. KESIMPULAN

Penelitian ini berhasil mengembangkan sistem *semantic search* untuk pencarian tafsir Al-Qur'an berbasis model IndoBERT yang dipadukan dengan FAISS sebagai mekanisme pengindeksan vektor. Melalui tahapan *preprocessing*, *fine-tuning*, pembentukan *embedding*, dan proses temu kembali berbasis *cosine similarity*, sistem mampu memahami hubungan semantik antara pertanyaan pengguna dan dokumen tafsir sehingga menghasilkan pencarian yang lebih relevan dibandingkan pendekatan berbasis pencocokan kata kunci. Hasil evaluasi menunjukkan bahwa sistem memperoleh nilai mAP sebesar 0,2128, MRR sebesar 0,2681, dan nDCG sebesar 0,2792 pada skenario Top-10, serta meningkat menjadi mAP 0,2301, MRR 0,2803, dan nDCG 0,3481 pada skenario Top-50. Hasil tersebut menunjukkan bahwa pendekatan yang diusulkan memiliki performa yang memadai dalam melakukan pencarian semantik pada domain tafsir Al-Qur'an. Selain itu, implementasi sistem dalam bentuk website memungkinkan pengguna memperoleh ayat, terjemahan, dan ringkasan Tafsir Al-Misbah melalui pertanyaan menggunakan bahasa alami.

Meskipun demikian, penelitian ini masih memiliki beberapa keterbatasan. Dataset yang digunakan hanya berasal dari Tafsir Al-Misbah sehingga cakupan pengetahuan sistem masih terbatas pada satu sumber tafsir. Selain itu, penelitian ini belum membandingkan performa metode yang diusulkan dengan metode *lexical retrieval* maupun *semantic retrieval* lainnya, sehingga keunggulan model belum dapat dibandingkan secara langsung dengan pendekatan lain. Oleh karena itu, penelitian selanjutnya dapat mengembangkan sistem dengan mengintegrasikan berbagai sumber tafsir Al-Qur'an, memperluas korpus data, melakukan perbandingan dengan metode lain seperti BM25, serta melibatkan evaluasi pengguna untuk memperoleh gambaran yang lebih komprehensif mengenai kualitas dan kegunaan sistem.

Penelitian selanjutnya dapat memperluas cakupan dataset dengan mengintegrasikan lebih banyak sumber tafsir Al-Qur'an serta mengeksplorasi model representasi bahasa yang lebih mutakhir untuk meningkatkan kualitas representasi semantik. Selain itu, evaluasi dapat diperluas melalui perbandingan dengan metode *semantic search* lainnya serta pengujian pada skenario penggunaan yang lebih beragam agar kinerja sistem dapat dinilai secara lebih komprehensif.

DAFTAR PUSTAKA

- [1] I. F. Hasanah, U. Hasanah, M. Makhzuniyah, dan A. N. K. B. Zawawi, "Qur' Anic Learning In The Digital Era : A Study On Digital Applications And Their Impact," *J. PAMATOR*, vol. 18, no. 1, hal. 146–159, 2025, doi: <https://doi.org/10.21107/pamator.v18i1.28172> Manuscript.
- [2] M. Huzaifa, B. Aqil, M. A. Haq, dan W. Zaghouni, *Arabic natural language processing for Qur'anic research : a systematic review*, vol. 56, no. 7. Springer Netherlands, 2023. doi: 10.1007/s10462-022-10313-2.
- [3] B. V. Kartika, M. J. Alfredo, dan G. P. Kusuma, "Fine-Tuned IndoBERT Based Model and Data Augmentation for Indonesian Language Paraphrase Identification," *Int. Inf. Eng. Technol. Assoc.*, vol. 37, no. 3, hal. 733–743, 2023, doi: <https://doi.org/10.18280/ria.370322> Received:
- [4] L. Trisnawati *dkk.*, "A proposed semantic keywords search engine for Indonesian Qur' an translation based on word embedding," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 35, no. 2, hal. 987–995, 2024, doi: 10.11591/ijeecs.v35.i2.pp987-995.
- [5] A. Malik, A. P. Gefadri, E. Sidik, dan A. P. Syadrina, "SoulScripture : Chatbot using Bidirectional Encoder Representations from Transformers as a Medium of Spiritual Guidance," *Khazanah J. Relig. Technol.*, vol. 2, no. 1, hal. 23–27, 2024, doi: <https://doi.org/10.15575/kjrt.v2i1.822> SoulScripture:
- [6] A. P. Putra, D. Purnami, S. Putri, A. A. Kt, dan A. Cahyawan, "Scientific Paper Recommendation System : Application of Sentence Transformers and Cosine Similarity Using arXiv Data," *J. Appl. Informatics Comput.*, vol. 9, no. 4, hal. 1374–1382, 2025, [Daring]. Tersedia pada: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [7] P. Chapman *dkk.*, "CRISP-DM 1.0: Step-by-Step Data Mining Guide," 2000.
- [8] K. Kowsari, K. J. Meimandi, M. Heidarysafa, dan S. Mendu, "Text Classification Algorithms : A Survey," *MDPI Inf.*, vol. 10, no. 1, hal. 1–68, 2020, doi: 10.3390/info10040150.
- [9] M. F. Ikhsan dan S. 'Uyun, "Analisis Perbandingan Metode Information Retrieval (IR) Pada Implementasi Pencarian Ayat Al-Qur' an," *J. SHIFT*, vol. 6, no. 1, 2026, doi: <https://doi.org/10.24252/shift.v6i1.272>.
- [10] I. Alpha dan E. Yulianti, "Academic expert finding using BERT pre-trained language model," *Int. J. Adv. Intell. Informatics*, vol. 10, no. 2, hal. 280–295, 2024, doi: <https://doi.org/10.26555/ijain.v10i2.1497>.
- [11] M. Pan, J. Wang, J. X. Huang, A. J. Huang, Q. Chen, dan J. Chen, "A probabilistic framework for integrating sentence-level semantics via BERT into pseudo-relevance feedback," *Inf. Process. Manag.*, vol. 59, no. 1, hal. 102734, 2022, doi: <https://doi.org/10.1016/j.ipm.2021.102734>.
- [12] H. Vo, Q. Nguyen, dan D. T. Tran, "Building an Information Retrieval System for University Documents Based on Generative AI Technologies," *VNUHCM J. Econ. Bus. Law*, vol. 10, no. 2,

- hal. 6508–6514, 2026, doi: <https://doi.org/10.32508/vnuhcmjebl.v10i2.1597>.
- [13] M. Alqarni, “Embedding Search for Quranic Texts based on Large Language Models,” *Int. Arab J. Inf. Technol.*, vol. 21, no. 2, hal. 243–256, 2024, doi: <https://doi.org/10.34028/iajit/21/2/71>.
- [14] N. E. Vijayan, “Enhancing Chatbot Response Relevance through Semantic Similarity Measures,” *J. Artif. Intell. Cloud Comput.*, vol. 1, no. 1, hal. 1–5, 2022, doi: [doi.org/10.47363/JAICC/2022\(1\)E182](https://doi.org/10.47363/JAICC/2022(1)E182).
- [15] T. Formal, B. Piwowarski, C. Lassance, dan S. Clinchant, “SPLADE v2: Sparse Lexical and Expansion Model for Information Retrieval,” *Proc. ACM Conf.*, vol. 1, no. 1, hal. 1–6, 2021, doi: <https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>.
- [16] M. Kamil dan D. Cakir, “Advances in Transformer-Based Semantic Search: Techniques, Benchmarks, and Future Directions,” *Turk. J. Math. Comput. Sci.*, vol. 17, no. 1, hal. 145–166, 2025, doi: [10.47000/tjmcs.1633092](https://doi.org/10.47000/tjmcs.1633092).