



Analisis Perbandingan Algoritma TF-IDF dan Word2Vec dalam Rekomendasi Film Berbasis Konten Letterboxd (Systematic Literature Review)

Diva Bulan^{1*}, Yulia Fatma²

¹⁻²Universitas Muhammadiyah Riau, Indonesia

Email: 220401040@student.umri.ac.id¹, yuliafatma@umri.ac.id²

*Penulis Korespondensi: 220401040@student.umri.ac.id

Abstract. The exponential growth of film streaming platforms has created significant challenges for users in discovering content that aligns with their preferences. Recommendation systems have proven to be a strategic solution. This study aims to compare the performance of Term Frequency-Inverse Document Frequency (TF-IDF) and Word2Vec algorithms in the context of content-based filtering recommendation systems, utilizing data from the Letterboxd platform. The research was conducted through a Systematic Literature Review (SLR) of 15 prior studies, with a selection process comprising the definition of inclusion and exclusion criteria, literature searches across academic databases, and comparative synthesis of findings. The review findings indicate that TF-IDF excels in weighting explicit terms within structured metadata such as genre, director, and cast, and can be implemented without a complex training process. Word2Vec, on the other hand, demonstrates superiority in capturing semantic relationships between words through dense vector representations, yet requires a large corpus and careful hyperparameter tuning. Based on the literature synthesis, TF-IDF combined with Cosine Similarity is deemed more optimal for a content-based film recommendation system using Letterboxd metadata, given the structured and explicit nature of the data. Recall@5 of 73% and Recall@10 of 80% reported in prior studies further support this conclusion

Keywords: Content-Based Filtering; Letterboxd; Natural Language Processing; Recommendation Film System; TF-IDF; Word2Vec.

Abstrak. Pertumbuhan eksponensial platform streaming film menciptakan tantangan bagi pengguna dalam menemukan konten yang sesuai preferensi mereka. Sistem rekomendasi telah terbukti menjadi solusi strategis, sebagaimana ditunjukkan oleh kontribusinya sebesar 35% terhadap pendapatan Amazon dan nilai USD 33,7 miliar bagi Netflix. Penelitian ini bertujuan membandingkan kinerja algoritma Term Frequency-Inverse Document Frequency (TF-IDF) dan Word2Vec dalam konteks sistem rekomendasi film berbasis kemiripan konten (*content-based filtering*) dengan memanfaatkan data dari platform Letterboxd. Penelitian dilakukan melalui Systematic Literature Review (SLR) terhadap 15 studi terdahulu, dengan tahapan seleksi meliputi penentuan kriteria inklusi dan eksklusi, pencarian literatur pada basis data akademik, serta sintesis temuan secara komparatif. Hasil kajian menunjukkan bahwa TF-IDF unggul dalam pembobotan term eksplisit pada metadata terstruktur seperti genre, sutradara, dan aktor, serta mudah diimplementasikan tanpa proses pelatihan yang kompleks. Adapun Word2Vec memiliki keunggulan dalam menangkap hubungan semantik antarkata melalui representasi vektor yang padat (*dense vector*), namun membutuhkan korpus data yang besar dan penyetelan hiperparameter yang cermat. Berdasarkan sintesis literatur, TF-IDF dengan metode Cosine Similarity dinilai lebih optimal untuk sistem rekomendasi film berbasis metadata Letterboxd, mengingat karakteristik data yang terstruktur dan eksplisit. Nilai Recall@5 sebesar 73% dan Recall@10 sebesar 80% yang ditemukan pada studi-studi terdahulu turut mendukung kesimpulan tersebut.

Kata Kunci: Letterboxd; Pemrosesan Bahasa Alami; Penyaringan Berbasis Konten; Sistem Rekomendasi Film; TF-IDF; Word2Vec.

1. LATAR BELAKANG

Industri hiburan digital, khususnya platform streaming film, telah mengalami pertumbuhan eksponensial dalam beberapa tahun terakhir. Pasar global layanan video on demand (VoD) mencapai nilai yang signifikan dengan proyeksi pertumbuhan berkelanjutan hingga tahun 2030. Di Indonesia sendiri, pendapatan dari langganan video on demand mencapai USD 411 juta pada tahun 2021 dengan penetrasi pengguna sebesar 16% dan diperkirakan akan terus meningkat menjadi 20% pada tahun 2025. Pertumbuhan ini dipercepat oleh pandemi COVID-19 yang mengubah pola konsumsi hiburan masyarakat dari bioskop konvensional ke platform digital. Industri perfilman Indonesia juga mencatat perkembangan positif dengan total 177 judul film yang diproduksi pada tahun 2023, meningkat dari 140 judul di tahun 2022 (Badan Pusat Statistik, 2024).

Seiring dengan meningkatnya jumlah konten film yang tersedia, pengguna platform streaming menghadapi tantangan besar dalam menemukan film yang sesuai dengan preferensi mereka. Amazon melaporkan bahwa 35% dari pendapatannya berasal dari sistem rekomendasi, sementara Netflix mengaitkan pendapatan sekitar USD 33,7 miliar dan kesuksesan retensi pelanggannya secara signifikan pada sistem rekomendasi yang mereka implementasikan. Kebutuhan akan informasi didorong oleh sejumlah besar penelitian dan teknologi.

2. KAJIAN TEORITIS

Mayoritas penelitian menggunakan dataset konvensional seperti MovieLens atau IMDb yang memiliki struktur metadata terbatas dan tidak mencerminkan dinamika platform komunitas modern seperti Letterboxd yang kaya akan user-generated content dan metadata kustom seperti tags, lists, dan detailed reviews. Kemudahan penggunaan website mengakibatkan bertambahnya dokumen teks yang berupa pendapat dan opini. Teknik yang berkembang saat penggalan dokumen teks waktu ini adalah text mining. Ada beberapa algoritma atau metode untuk analisis sentimen, antara lain term frequency-inverse document frequency dan word2vec. Di dalam penelitian ini akan menganalisa kinerja dua algoritma natural language processing, yakni TF-IDF dan Word2Vec.

3. METODE PENELITIAN

Penelitian ini menggunakan metode Systematic Literature Review (SLR) sebagai pendekatan utama dalam mengumpulkan dan menganalisis data dengan menggunakan framework Kitchenham & Charters.

Proses penelitian dilakukan secara sistematis melalui beberapa tahapan, yaitu perencanaan, pencarian literatur, seleksi sumber, ekstraksi data, dan analisis sintesis yang berasal dari Arxiv dan Google Scholar. Mengingat kompleksitas permasalahan yang dikaji yang melibatkan evaluasi komparatif dua pendekatan algoritmik berbeda dalam domain Natural Language Processing (NLP) dan Information Retrieval—diperlukan pendekatan metodologis yang mampu mengintegrasikan berbagai temuan empiris, teoretis, dan praktis dari penelitian-penelitian terdahulu. Sistem rekomendasi berbasis konten (content-based filtering) yang memanfaatkan data tekstual dari platform Letterboxd memerlukan pemahaman mendalam tentang bagaimana representasi teks mempengaruhi akurasi dan relevansi rekomendasi yang dihasilkan sehingga pemilihan artikel berdasarkan kesesuaian rekomendasi konten berbasis teks.

Pertanyaan Penelitian:

- a) RQ 1: Algoritma apa yang paling dominan dalam sistem rekomendasi film berbasis konten?
- b) RQ 2: Tantangan penelitian apa yang terjadi dalam kasus rekomendasi film berbasis konten?
- c) RQ 3: Peluang penelitian apa yang ada dan dapat dilakukan di masa yang akan datang terhadap rekomendasi berbasis konten?
- d) RQ 4: Apa yang membedakan algoritma TF-IDF dan Word2Vec dalam rekomendasi berbasis konten?

4. HASIL DAN PEMBAHASAN

Berdasarkan hasil seleksi literatur, sebanyak 15 artikel dipilih sebagai objek analisis utama. Seluruh artikel membahas implementasi algoritma TF-IDF dan Word2Vec dalam konteks rekomendasi.

Tabel 1. Perbandingan Performa Metode *Deep Learning* Dalam Algoritma Konteks Rekomendasi.

No.	Judul Artikel	Metode	Jumlah Dataset	Akurasi Word2Vec	Akurasi TF-IDF	Kelebihan	Kekurangan
1.	Enhancing mortality prediction in cardiac arrest ICU patients through meta-modeling of structured clinical data from MIMIC-	TF-IDF dan Word2Vec	458.000	92%	75%	TF-IDF mampu mengidentifikasi term-term penting yang kontekstual dengan memberikan bobot tinggi pada kata yang sering muncul	Metode ini juga tidak dapat menangkap hubungan semantik atau kontekstual antar kata - TF-IDF hanya mempertimbangkan frekuensi

	IV (Mamatov and Kellmeyer 2025)					dalam dokumen tertentu namun jarang di seluruh corpus.	kemunculan kata tanpa memahami makna atau hubungan antar term.
2.	Predicting person-level injury severity using crash narratives: A balanced approach with roadway classification and natural language process techniques (Majidi et al. 2025)	Word2Vec dan TF-IDF	3.200	70%	75%	kemampuan TF-IDF untuk memberikan inherent feature selection dengan secara otomatis down-weighting common atau non-discriminative terms.	TF-IDF tidak dapat menangkap semantic relationships atau contextual meaning antar kata dalam crash narratives - metode ini hanya menghitung term frequency tanpa memahami makna atau hubungan konseptual.
3.	Multi-Label Clinical Text Eligibility Classification and Summarization System (Yerramsetty and Fathimah 2025)	TF-IDF	288		66%	TF-IDF efektif sebagai fitur dasar untuk teks klinis karena dapat menonjolkan istilah medis penting dan mengurangi dominasi kata umum.	Kurang mampu menangkap relasi semantik antar istilah medis, sehingga informasi konteks kompleks dalam catatan klinis tidak terwakili penuh.
4.	Irony Detection in Urdu Text: A Comparative Study Using Machine Learning Models and Large Language Models (Ahmad et al. n.d.)	Word2Vec	1.950	89.18%		Mampu menangkap informasi semantik global melalui averaging word embeddings	Kurang efektif untuk fenomena linguistik yang subtle dan memerlukan pemahaman konteks mendalam
5.	Type and Complexity Signals in Multilingual Question Representations (Kokot and Poelman 2025)	TF-IDF	9000		80.9 %	Dapat menangkap pola frekuensi permukaan yang reliable untuk klasifikasi tipe pertanyaan	Kurang efektif untuk bahasa yang memerlukan integrasi kontekstual (Jepang, Korea, Inggris, Finlandia) dengan selisih performa 10-20% dibanding model neural
6.	Automated Research	TF-IDF	1.200.000		69%	Teknik yang mature dan teruji	Hanya bekerja pada level kata,

	Article Classification and Recommendation Using NLP and Machine Learning (Rahman et al. 2025)					untuk berbagai tugas termasuk klasifikasi teks, clustering, dan information retrieval	tidak mempertimbangkan urutan atau konteks kalimat
7.	Graph-Structured Data Analysis of Component Failure in Autonomous Cargo Ships Based on Feature Fusion (Zhang et al. 2024)	Word2Vec	6.150	73.5%		Dengan dimensi 100, dapat menyeimbangkan kemampuan representasi semantik dan kompleksitas komputasi	Tidak cocok untuk teks yang lebih panjang seperti Failure Mode dan Failure Reason yang berupa frasa atau kalimat pendek (penelitian ini harus menggunakan BERT untuk kasus tersebut)
8.	Prediction of Coffee Ratings Based On Influential Attributes Using SelectKBest and Optimal Hyperparameters (Agyemang et al. 2025)	TF-IDF			94%	Memberikan bobot lebih tinggi pada kata-kata yang penting dalam dokumen individual tetapi tidak umum di seluruh korpus	TF-IDF menghasilkan representasi sparse yang memerlukan proses seleksi fitur tambahan menggunakan SelectKBest
9.	Analyzing the Evolution of Graphs and Texts (Analyzing the Evolution of Graphs and Texts 2023)	Word2Vec	41.960	85%		Efektif dalam merepresentasikan kata-kata sebagai vektor dense (60-300 dimensi)	Memerlukan dua tahap: pertama menghitung co-occurrence matrix yang besar dan sparse, kemudian melakukan self-supervised pre-training neural network
10.	Analisis Sentimen Review Film Menggunakan TF-IDF dan Support Vector Machine (Gifari et al. 2022)	TF-IDF dan SVM	200		85%	Penelitian ini membuktikan bahwa TF-IDF dapat menghasilkan representasi numerik yang baik untuk dokumen teks, dengan sistem mencapai akurasi 85%, precision 100%, recall	TF-IDF mengalami kesulitan menangani kata-kata tidak baku dan singkatan yang umum di Twitter, yang menyebabkan hilangnya informasi penting setelah proses preprocessing.

11	Model Klasifikasi Berita Palsu Menggunakan <i>Bidirectional LSTM</i> Dan <i>Word2Vec</i> Sebagai Vektorisasi (Amalia et al. 2022)	Word2Vec		77,3%	70%, dan F1-Score 82%. Penelitian ini menunjukkan bahwa <i>Word2Vec</i> memiliki beberapa keunggulan signifikan dalam vektorisasi teks. Metode ini mampu memberikan representasi vektor yang padat dan pendek sehingga tidak memakan banyak memori dan sumber daya komputasi.	Meskipun efektif, <i>Word2Vec</i> memerlukan pengaturan hyperparameter yang tepat seperti embedding size dan window size untuk menghasilkan performa optimal.
12	Sistem Rekomendasi Content-based Filtering Menggunakan TF-IDF Vector Similarity Untuk Rekomendasi Artikel Berita (Huda, Fajarudin, and Hadinegoro 2022)	TF-IDF	345		80% Penelitian ini mendemonstrasikan keefektifan TF-IDF dalam sistem rekomendasi artikel berita. TF-IDF mampu mengidentifikasi kata-kata penting dalam dokumen dengan memberikan skor tinggi pada kata yang sering muncul dalam dokumen tertentu namun jarang muncul di dokumen lain.	Hal ini menunjukkan bahwa TF-IDF kurang efektif dalam membedakan antara konten substantif dan elemen teknis dalam dokumen. Selain itu, performa sistem sangat bergantung pada kualitas dan kuantitas dataset
13	Feature Expansion Word2Vec for Sentiment Analysis of Public Policy in Twitter (Royyan and Setiawan 2022)	Word2Vec	16.597	78.99%	<i>Word2Vec</i> mampu mengelompokkan kata-kata yang memiliki kesamaan semantik ke dalam satu ruang vektor, kemudian membentuk corpus similarity yang berisi kumpulan kata-kata serupa.	Ukuran corpus similarity sangat mempengaruhi hasil - corpus yang terlalu kecil (Tweet corpus: 10,765 kata) menghasilkan performa yang jauh lebih buruk dibanding corpus yang lebih besar (IndoNews: 178,568 kata atau IndoNews+Tweet: 179,636 kata).

14	Implementasi Metode Term Frequency-Inverse Document Frequency dan Cosine Similarity pada Bot Telegram untuk Mendukung Rekomendasi Laptop Berdasarkan Preferensi Pengguna (Safira, Anwar, and Aldrin 2025)	TF-IDF	1303	80%	Metode ini mampu memberikan bobot yang tepat pada hubungan kata terhadap dokumen, sehingga chatbot dapat mengidentifikasi spesifikasi laptop yang sesuai dengan preferensi pengguna.	TF-IDF mengalami kesulitan ketika berhadapan dengan pertanyaan yang berada di luar konteks penggunaannya atau pertanyaan dengan struktur yang tidak umum.
15	The Accuracy Comparison Between Word2Vec and FastText On Sentiment Analysis of Hotel Reviews (Khomsah, Ramadhani, and Wijayanto 2022)	Word2Vec	3145	93%	Word2Vec menunjukkan performa yang konsisten ketika dikombinasikan dengan ensemble learning. Dengan menggunakan arsitektur Skip-gram, Word2Vec mampu menangkap konteks kata berdasarkan kata-kata di sekitarnya melalui mekanisme window size.	Penelitian menunjukkan bahwa Word2Vec bekerja sangat baik untuk sentimen positif (precision 97%, recall 93%) namun kurang optimal untuk sentimen negatif (precision 75%, recall 87%).

Perbandingan Performa Berdasarkan Algoritma

Berdasarkan hasil dari 15 jurnal yang dikaji, terdapat beberapa rentang hasil akurasi di antara dua algoritma, yakni sebagai berikut :

Tabel 2. Hasil Perbandingan Algoritma TF-IDF dan Word2Vec.

Algoritma	Jumlah Penelitian	Rentang Akurasi
TF-IDF	7	70% - 93%
Word2Vec	8	66% - 94%

Berdasarkan Tabel 2, TF-IDF menunjukkan rentang akurasi yang relatif tinggi, yaitu 70% – 93%. Namun, nilai akurasi yang sangat tinggi ini perlu ditafsirkan secara hati-hati karena sebagian besar diperoleh dari dataset berukuran kecil hingga menengah yang berpotensi menimbulkan bias dan risiko overfitting. Sebaliknya, Word2Vec menunjukkan performa yang lebih rendah dengan rentang akurasi 66% – 94%, khususnya pada dataset berukuran menengah hingga besar.

Tantangan Penelitian yang Dihadapi

Meskipun potensinya sangat besar, penelitian yang secara spesifik mengeksplorasi sistem rekomendasi film berbasis konten dengan memanfaatkan dataset Letterboxd masih sangat terbatas. Minimnya eksplorasi ini mengindikasikan adanya celah penelitian yang signifikan, terutama dalam memahami bagaimana fitur-fitur khas platform ini, seperti metadata kustom, tag komunitas, dan ulasan pengguna (*user-generated content*) dapat diintegrasikan secara efektif ke dalam pipeline sistem rekomendasi. Tantangan teknis yang tidak kalah krusial adalah kompleksitas linguistik yang terkandung dalam teks ulasan pengguna Letterboxd. Komentar pada platform ini kerap mengandung ironi, sarkasme, dan rujukan sinematik yang bersifat kontekstual, sehingga pendekatan berbasis frekuensi kata seperti TF-IDF berpotensi gagal menangkap makna semantik yang sesungguhnya. Kondisi ini menuntut penerapan teknik NLP yang lebih canggih, seperti analisis sentimen berbasis leksikon atau model berbasis Transformer, guna merepresentasikan teks ulasan secara lebih akurat. Selain itu, rating numerik dari pengguna yang beragam latar belakang budaya dan preferensinya dapat menghasilkan data yang bersifat subjektif dan tidak konsisten, sehingga diperlukan strategi normalisasi atau pendekatan *sentiment-aware recommendation* untuk meningkatkan kualitas representasi fitur dalam model rekomendasi.

Peluang Penelitian di Masa Depan

Berdasarkan hasil analisis komparatif terhadap penelitian terdahulu, dapat disimpulkan bahwa belum terdapat algoritma yang sepenuhnya optimal untuk sistem rekomendasi film berbasis konten pada dataset Letterboxd. Sebagian besar penelitian masih berfokus pada peningkatan akurasi semata, sementara aspek generalisasi model terhadap data yang beragam, ketahanan terhadap *cold-start problem*, serta efisiensi komputasi pada skala besar belum dikaji secara mendalam. Selain itu, pendekatan yang mengintegrasikan analisis sentimen dari ulasan pengguna (*sentiment-aware recommendation*), arsitektur rekomendasi hibrida (*hybrid recommender systems*) yang menggabungkan *content-based* dan *collaborative filtering*, serta pemanfaatan model berbasis Transformer atau Large Language Model (LLM) untuk pemahaman konteks semantik yang lebih kaya, masih terbuka luas sebagai arah pengembangan penelitian selanjutnya. Oleh karena itu, penelitian ini dirancang untuk mengkaji perbandingan algoritma TF-IDF dan Word2Vec secara lebih komprehensif, dengan mempertimbangkan keseimbangan antara performa akurasi, efisiensi komputasi, dan potensi penerapannya dalam sistem rekomendasi berbasis teks nyata.

Perbedaan Kedua Algoritma

Berdasarkan hasil evaluasi komparatif terhadap literatur terdahulu, kedua algoritma memiliki karakteristik teknis yang berbeda secara fundamental dalam merepresentasikan teks. TF-IDF merupakan metode pembobotan statistik yang mengukur relevansi sebuah kata dalam dokumen secara relatif terhadap seluruh koleksi korpus. Secara matematis, nilai TF-IDF dihitung melalui perkalian dua komponen: *Term Frequency* (TF), yaitu rasio frekuensi kemunculan kata dalam dokumen, dan *Inverse Document Frequency* (IDF), yaitu logaritma invers dari proporsi dokumen yang mengandung kata tersebut. Pendekatan ini menghasilkan representasi vektor yang bersifat *sparse* dan berdimensi tinggi, di mana setiap dimensi merepresentasikan satu kata unik dalam kosakata. Dari sisi kompleksitas komputasi, TF-IDF tergolong efisien dengan kompleksitas $O(n \times m)$, di mana n adalah jumlah dokumen dan m adalah ukuran kosakata, sehingga tidak memerlukan proses pelatihan (*training-free*) dan dapat dijalankan pada sumber daya komputasi yang terbatas.

Word2Vec, sebaliknya, merupakan model jaringan saraf tiruan (*neural network*) yang melatih representasi kata dalam bentuk vektor padat (*dense vector*) berdimensi rendah, umumnya antara 100 hingga 300 dimensi, dengan cara menangkap hubungan kontekstual dan semantik antarkata. Model ini tersedia dalam dua arsitektur: *Continuous Bag-of-Words* (CBOW) yang memprediksi sebuah kata berdasarkan konteks kata-kata di sekitarnya, dan *Skip-gram* yang melakukan sebaliknya, yakni memprediksi konteks dari sebuah kata target.

Berbeda dengan TF-IDF, Word2Vec mampu menangkap kemiripan semantik secara implisit, sehingga kata-kata yang berkaitan makna akan memiliki jarak vektor yang berdekatan dalam ruang representasi. Namun demikian, kompleksitas pelatihannya jauh lebih tinggi, yakni $O(V \times D \times E)$ di mana V adalah ukuran kosakata, D adalah dimensi vektor, dan E adalah jumlah epoch, sehingga membutuhkan korpus data yang besar dan sumber daya komputasi yang memadai untuk menghasilkan representasi yang optimal.

Perbedaan mendasar antara keduanya dapat dirangkum sebagai berikut: TF-IDF unggul dalam efisiensi dan interpretabilitas untuk metadata terstruktur, sementara Word2Vec lebih unggul dalam pemahaman semantik kontekstual namun dengan biaya komputasi yang lebih besar.

5. KESIMPULAN DAN SARAN

Berdasarkan sintesis komparatif terhadap 15 penelitian terdahulu, TF-IDF dan Word2Vec menunjukkan profil kinerja yang berbeda dan saling melengkapi. TF-IDF terbukti efektif dalam mengidentifikasi term-term penting pada metadata terstruktur, mudah diimplementasikan tanpa proses pelatihan, dan konsisten menghasilkan performa tinggi pada tugas-tugas berbasis kata kunci eksplisit. Word2Vec, di sisi lain, unggul dalam menangkap hubungan semantik antarkata melalui representasi dense vector, namun sangat bergantung pada ketersediaan korpus berukuran besar dan konfigurasi hiperparameter yang tepat.

Berdasarkan karakteristik tersebut, penelitian ini menyimpulkan bahwa TF-IDF dengan Cosine Similarity lebih sesuai untuk sistem rekomendasi film berbasis metadata Letterboxd, mengingat data yang tersedia bersifat terstruktur dan eksplisit — seperti genre, sutradara, aktor, dan deskripsi plot. Nilai Recall@5 sebesar 73% dan Recall@10 sebesar 80% yang ditemukan pada studi-studi terdahulu turut mendukung kesimpulan ini, meskipun perlu ditegaskan bahwa angka tersebut berasal dari eksperimen penelitian sebelumnya dan bukan dari pengujian langsung dalam penelitian ini.

Penelitian ini memiliki beberapa keterbatasan yang perlu diakui secara eksplisit. Pertama, kesimpulan yang dihasilkan sepenuhnya didasarkan pada sintesis literatur melalui pendekatan SLR, sehingga tidak dapat menggantikan validasi empiris melalui eksperimen langsung pada dataset Letterboxd. Kedua, perbandingan kinerja kedua algoritma dilakukan secara tidak langsung, karena masing-masing diuji pada dataset dan konteks yang berbeda-beda dalam literatur yang dikaji.

Ketiga, penelitian ini belum mempertimbangkan variabel kontekstual seperti distribusi bahasa ulasan, kelengkapan metadata, dan dinamika preferensi pengguna yang dapat memengaruhi performa algoritma secara signifikan.

Untuk penelitian selanjutnya, pendekatan hybrid recommender system yang mengintegrasikan TF-IDF untuk pemrosesan metadata terstruktur dengan Word2Vec atau model berbasis Transformer untuk analisis ulasan naratif berpotensi menghasilkan sistem rekomendasi yang lebih akurat dan kontekstual. Eksplorasi terhadap sentiment-aware recommendation berbasis ulasan pengguna Letterboxd juga dinilai menjanjikan sebagai arah pengembangan yang layak diuji secara empiris..

DAFTAR REFERENSI

- Agyemang, Edmund, Lawrence Agbota, Vincent Agbenyeavu, Peggy Akabuah, Bismark Bimpong, and Christopher Attafuah. 2025. "Prediction of Coffee Ratings Based On Influential Attributes Using SelectKBest and Optimal Hyperparameters." 1–13. <http://arxiv.org/abs/2509.18124>.
- Ahmad, Fiaz, Nisar Hussain, Amna Qasim, Momina Hafeez, Muhammad Usman Grigori, and Alexander Gelbukh. n.d. "Irony Detection in Urdu Text : A Comparative Study Using Machine Learning Models and Large Language Models." 1–5.
- Amalia, Junita, Juanda Pakpahan, Melani Pakpahan, and Yeni Panjaitan. 2022. "Model Klasifikasi Berita Palsu Menggunakan Bidirectional LSTM Dan Word2Vec Sebagai Vektorisasi." 9(4):3319–31.
- Analyzing the Evolution of Graphs and Texts. 2023. (May).
- Gifari, Okta Ihza, Muh. Adha, Fernandito Freddy, and Fernandito Freddy Setlight Durrand. 2022. "Analisis Sentimen Review Film Menggunakan TF-IDF Dan Support Vector Machine." *Journal of Information Technology* 2(1):36–40. doi:10.46229/jifotech.v2i1.330.
- Huda, Arif Akbarul, Rohmad Fajarudin, and Arifiyanto Hadinegoro. 2022. "Sistem Rekomendasi Content-Based Filtering Menggunakan TF-IDF Vector Similarity Untuk Rekomendasi Artikel Berita." *Building of Informatics, Technology and Science (BITS)* 4(3):1679–86. doi:10.47065/bits.v4i3.2511.
- Khomsah, Siti, Rima Dias Ramadhani, and Sena Wijayanto. 2022. "The Accuracy Comparison Between Word2Vec and FastText On." *Rekayasa Sistem Dan Teknologi Informasi* 5(158):352–58.
- Kokot, Robin, and Wessel Poelman. 2025. "Type and Complexity Signals in Multilingual Question Representations." 411–25. doi:10.18653/v1/2025.mrl-main.28.
- Majidi, Mohammad Zana, Sajjad Karimi, Teng Wang, Robert Kluger, and Reginald Souleyrette. 2025. "Predicting Person-Level Injury Severity Using Crash Narratives: A Balanced Approach with Roadway Classification and Natural Language Process Techniques."
- Mamatov, Nursultan, and Philipp Kellmeyer. 2025. "Enhancing Mortality Prediction in Cardiac Arrest ICU Patients through Meta-Modeling of Structured Clinical Data from MIMIC-IV." <http://arxiv.org/abs/2510.18103>.
- Rahman, Shadikur, Hasibul Karim Shanto, Umme Ayman Koana, and Syed Muhammad Danish. 2025. "Automated Research Article Classification and Recommendation Using NLP and ML." <http://arxiv.org/abs/2510.05495>.
- Royyan, Alvi Rahmy, and Erwin Budi Setiawan. 2022. "Feature Expansion Word2Vec for Sentiment Analysis of Public Policy in Twitter." *Jurnal RESTI* 6(1):78–84. doi:10.29207/resti.v6i1.3525.
- Safira, Stevani Dean, Nanda Rosma Anwar, and Medistra Aldrin. 2025. "Implementasi Metode Term Frequency-Inverse Document Frequency Dan Cosine Similarity Pada Bot Telegram Untuk Mendukung Rekomendasi Laptop Berdasarkan Preferensi Pengguna Implementation of the TF-IDF and Cosine Similarity Methods in Telegram Bots." 4(1):63–73.

- Yerramsetty, Surya Tejaswi, and Almas Fathimah. 2025. "Multi-Label Clinical Text Eligibility Classification and Summarization System." <https://arxiv.org/pdf/2510.13115>.
- Zhang, Zizhao, Tianxiang Zhao, Yu Sun, Liping Sun, and Jichuan Kang. 2024. "Graph-Structured Data Analysis of Component Failure in Autonomous Cargo Ships Based on Feature Fusion." 1–22.