

Integrasi *Word Embeddings* & Web Scraping IMDb untuk Rekomendasi Film Berbasis Kata Kunci

Integrating Word Embeddings and IMDb Web Scraping for Keyword-Based Movie Recommendation

Andani Chacha Cahya Dewi¹, Deni Arifianto², Nanda Kurnia Wardati³

^{1,2,3} Program Studi Sistem Informasi, Fakultas Teknik, Universitas Muhammadiyah Jember

¹andanichachacahyadewi@gmail.com, ²deniarifianto@unmuhjember.ac.id,

³nandakurniawardati@unmuhjember.ac.id

Received: June 20, 2026 | Revised: July 18, 2026 | Accepted: July 31, 2026

Abstrak

Pertumbuhan pesat industri film dan platform streaming memicu information overload dan filter bubble yang menyulitkan pengguna menemukan tontonan sesuai preferensi alur cerita. Pendekatan content-based filtering berbasis frekuensi kata (TF-IDF) pada penelitian terdahulu memiliki keterbatasan dalam menangkap kesamaan makna antarkata (semantic gap), dan umumnya masih bergantung pada dataset publik tunggal serta input judul film acuan (seed movie). Penelitian ini menggabungkan dataset publik dengan hasil web scraping IMDb (populasi maksimal 5.000 judul film) dan menerapkan model Word2Vec arsitektur Skip-gram untuk merepresentasikan sinopsis film ke dalam vektor semantik berdimensi 200, dipadukan dengan Cosine similarity untuk menghitung kemiripan antara kata kunci (keyword) bebas dari pengguna dan sinopsis film tanpa memerlukan judul film acuan. Data dibagi menggunakan metode Holdout (80:20) dan performa algoritma dievaluasi pada jendela Top-3 Recommendation memakai metrik Precision@K, Recall@K, dan Mean Reciprocal Rank (MRR), dengan ground truth hasil validasi dua pakar melalui Inter-Annotator Agreement. Pengujian terhadap 25 kueri menghasilkan Precision@3 sebesar 0,5333, Recall@3 sebesar 0,7800, dan MRR sebesar 0,7300. Hasil ini menunjukkan bahwa integrasi word embeddings dan web scraping mampu menghasilkan rekomendasi film yang relevan secara semantik dari masukan kata kunci bebas, meskipun perbandingan dengan baseline masih diperlukan untuk klaim kinerja yang lebih definitif.

Kata kunci: sistem rekomendasi, *content-based filtering*, *word embeddings*, *Word2Vec*, *cosine similarity*

Abstract

The rapid growth of the film industry and streaming platforms has led to information overload and filter bubbles that make it difficult for users to find content matching their narrative preferences. Prior content-based filtering approaches relying on word-frequency methods (TF-IDF) suffer from a semantic gap and commonly depend on a single public dataset and a reference title (seed movie) as input. This study combines a public dataset with IMDb web-scraping results (a maximum population of 5,000 titles) and applies a Skip-gram Word2Vec model to represent movie synopses as 200-dimensional semantic vectors, paired with Cosine similarity to measure the closeness between a user's free-text keyword and movie synopses without requiring a seed movie. Data were split using an 80:20 Holdout method, and algorithm performance was evaluated on a Top-3 Recommendation window using Precision@K, Recall@K, and Mean Reciprocal Rank (MRR), with ground truth validated by two experts through Inter-Annotator Agreement. Testing on 25 queries produced a Precision@3 of 0.5333, Recall@3 of 0.7800, and MRR of 0.7300. These results indicate that integrating word embeddings with web scraping yields semantically relevant movie recommendations from free keyword input, though comparisons with baseline methods are needed for more definitive performance claims.

Keywords: *recommender system, content-based filtering, word embeddings, Word2Vec, cosine similarity*

1. PENDAHULUAN

Industri film dan layanan streaming digital tumbuh pesat, dengan ribuan judul baru dirilis setiap tahun dan katalog berisi puluhan ribu konten yang dapat diakses pengguna secara daring [1], [2]. Kondisi ini memperluas pilihan tontonan, tetapi juga memunculkan *information overload* yang menyulitkan pengguna menentukan film sesuai minat personal. Sistem rekomendasi berperan penting sebagai komponen strategis untuk menyaring dan mempersonalisasi konten pada platform digital [3].

Meskipun sistem rekomendasi telah banyak diterapkan, pengguna masih sering terjebak dalam *filter bubble*, yaitu rekomendasi yang cenderung berputar pada film populer atau genre yang berulang, sehingga menimbulkan *choice paralysis* ketika pengguna dihadapkan pada terlalu banyak pilihan umum [4]. Penelitian ini dirancang untuk memberi nilai tambah melalui pencarian berbasis kata kunci (*keyword*) bebas: dengan memasukkan kata kunci yang mendeskripsikan alur cerita yang diinginkan, sistem membantu pengguna menemukan *hidden gems* yang secara semantik relevan dengan kedalaman cerita yang disukai [5].

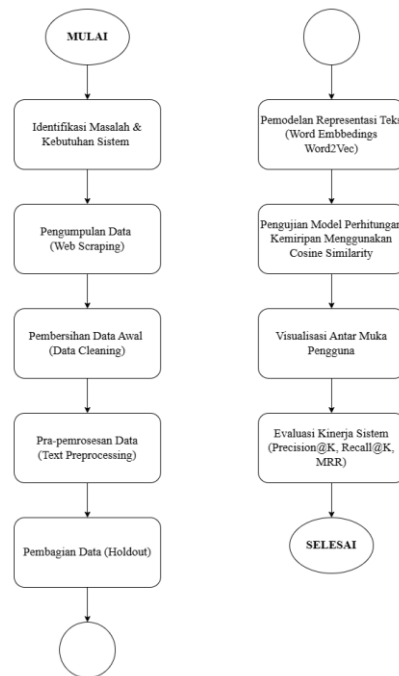
Tantangan lain pada *content-based filtering* adalah ketersediaan dan kekayaan teks deskripsi film. Dataset publik memberi kontribusi besar karena terstandarisasi dan andal [6], serta membantu meminimalkan kendala *cold start* pada film baru selama deskripsi kontennya memadai [7]. Namun agar representasi makna semantik lebih maksimal, korpus perlu diperkaya dengan volume dan variasi kata yang lebih besar, sehingga penelitian ini mengintegrasikan *web scraping* dari IMDb sebagai pelengkap dataset publik [8].

Pendekatan *content-based filtering* konvensional umumnya mengekstraksi teks berbasis frekuensi kata seperti TF-IDF, yang memiliki keterbatasan dalam memahami konteks makna naratif: kata-kata yang berbeda secara leksikal namun bermakna serupa dapat dianggap tidak berhubungan, sehingga kemiripan yang dihasilkan belum sepenuhnya mencerminkan kesamaan isi [9]. *Word embeddings* diperkenalkan sebagai alternatif untuk merepresentasikan teks ke ruang vektor sambil mempertahankan makna semantis dan hubungan antarkata [10], sehingga kata kunci pengguna dapat dipetakan lebih presisi terhadap vektor deskripsi film dibandingkan metode berbasis frekuensi kata, tanpa bergantung pada judul film acuan (*seed movie*).

Penelitian-penelitian terdahulu pada *content-based filtering* dan *word embeddings* untuk sistem rekomendasi umumnya masih berfokus pada pemanfaatan dataset publik sekunder secara tunggal dan bergantung pada judul film acuan dalam pencariannya [6], [9]. Beberapa studi terbaru telah mengeksplorasi penggunaan model *Transformer* seperti BERT untuk representasi semantik [13] dan pendekatan *hybrid* yang menggabungkan metadata film [14]; namun pendekatan tersebut umumnya menuntut sumber daya komputasi yang lebih besar dan tetap mengandalkan *seed movie* sebagai input. Berbeda dengan Kumar dan Singh [13] yang menerapkan BERT untuk representasi semantik sinopsis namun tetap bergantung pada *seed movie* dan dataset statis, serta Zhang et al. [14] yang mengintegrasikan metadata film dalam pendekatan *hybrid* namun tidak menangani input berbahasa Indonesia, penelitian ini mengisi celah tersebut melalui tiga kontribusi komplementer: (1) representasi semantik murni berbasis Word2Vec *Skip-gram* tanpa *seed movie*; (2) korpus dinamis hasil kombinasi dataset publik dan *web scraping* IMDb; dan (3) antarmuka pencarian berbasis terjemahan otomatis kata kunci Bahasa Indonesia yang belum ditemukan dalam literatur mutakhir [15]. Penelitian ini mengintegrasikan *word embeddings* (Word2Vec *Skip-gram*) dan *Cosine similarity* dalam sistem rekomendasi tanpa *seed movie*, dan mengukur kualitas urutan rekomendasi menggunakan metrik Top-N Recommendation (Precision@K, Recall@K, dan MRR) dengan ground truth yang divalidasi melalui penilaian pakar (*expert judgment*).

2. METODE PENELITIAN

Penelitian ini menerapkan pipeline data science yang terdiri atas pengumpulan data, pembersihan dan prapemrosesan teks, pembagian data (*Holdout*), pemodelan representasi teks (*Word2Vec*), pengujian kemiripan (*Cosine similarity*), dan evaluasi kinerja algoritma. Sistem kemudian diimplementasikan sebagai prototype berbasis web menggunakan *Flask* untuk menguji masukan pengguna dan menampilkan hasil rekomendasi. Alur penelitian ditunjukkan pada Gambar 1.



Gambar 1. Tahapan Penelitian

2.1 Pengumpulan Data

Data bersumber dari dataset publik *Internet Movie Database* (IMDb) yang diperkaya secara komplementer melalui *web scraping* berbasis Python (Selenium dan BeautifulSoup), untuk menjaga kesegaran dan relevansi kosakata sinopsis dengan perkembangan industri film terbaru [8]. Populasi data dibatasi maksimal 5.000 judul film berbahasa Inggris, dirilis hingga April 2026, dengan rating audiens minimal 6,0 dan wajib memiliki teks sinopsis lengkap; baris data yang kosong dieliminasi otomatis. Sinopsis dipertahankan dalam bahasa aslinya (Inggris) untuk menjaga kemurnian makna semantik saat diproses oleh algoritma pemrosesan bahasa alami.

2.2 Pembersihan dan Prapemrosesan Teks

Data mentah hasil scraping dibersihkan dari duplikasi, nilai kosong, dan noise, kemudian melalui tahap text preprocessing meliputi *case folding*, tokenisasi, penghapusan *stopword*, dan lemmatisasi menggunakan pustaka NLTK, guna menghasilkan korpus teks yang bersih dan siap dilatih [11].

2.3 Pembagian Data (Holdout)

Dataset dibagi menggunakan metode Holdout Split (*train_test_split* dari scikit-learn) dengan rasio 80:20, menghasilkan 4.000 data latih untuk melatih model Word2Vec dan 1.000 data uji yang dipisahkan secara ketat untuk pengujian objektif tahap Cosine similarity dan evaluasi metrik Top-N. Untuk memastikan reproduktibilitas, pembagian data dilakukan dengan *random_state=42* agar partisi yang sama dapat diperoleh kembali pada percobaan ulang.

2.4 Pemodelan Word Embeddings (Word2Vec)

Representasi vektor kata dilatih menggunakan model Word2Vec arsitektur Skip-gram melalui pustaka Gensim [10], yang dipilih karena lebih sensitif terhadap kata-kata yang jarang muncul dan mampu menangkap relasi semantik pada korpus naratif berukuran menengah. Pemilihan arsitektur Skip-gram (sg=1) didasarkan pada karakteristik korpus sinopsis film yang mengandung banyak kata naratif dengan frekuensi rendah namun bermakna tinggi; Skip-gram terbukti lebih unggul dari CBOW pada kondisi ini [10]. Konfigurasi parameter pelatihan ditunjukkan pada Tabel 1.

Tabel 1. Konfigurasi Parameter Word2Vec (Skip-gram)

Parameter	Nilai	Keterangan
sg	1	Arsitektur Skip-gram
vector_size	200	Dimensi vektor kata; keseimbangan antara kapasitas representasi semantik dan efisiensi memori pada dataset 5.000 dokumen.
window	8	Jarak kata target-konteks; lebih besar dari default (5) untuk mengakomodasi klausa sinopsis film yang panjang.
min_count	2	Ambang frekuensi kata; memastikan kata sangat jarang (noise/typo) tidak membentuk ruang vektor.
epochs	30	Iterasi pelatihan; cukup untuk konvergensi pada korpus berukuran menengah.
seed	42	Nilai seed untuk reproduktibilitas hasil pelatihan.

Seluruh proses pelatihan dan pengujian dijalankan pada lingkungan Python 3.11 dengan pustaka utama Gensim 4.3.2, scikit-learn 1.4.0, NLTK 3.8.1, Flask 3.0.2, Selenium 4.18.1, dan BeautifulSoup4 4.12.3, pada mesin dengan RAM 16 GB dan prosesor Intel Core i5 generasi ke-12. Tidak ada GPU yang digunakan dalam pelatihan model ini.

Penelitian ini tidak menyertakan eksperimen komparasi langsung antara Word2Vec dan baseline TF-IDF maupun model embedding alternatif (FastText, GloVe) dalam satu skenario uji terkontrol. Keputusan ini diambil secara metodologis karena fokus penelitian adalah validasi pipeline end-to-end yang mencakup web scraping, prapemrosesan teks, dan evaluasi berbasis expert judgment penambahan varian komparasi akan memperluas ruang lingkup melampaui sumber daya yang tersedia dan mengurangi kedalaman analisis pada setiap komponen. Perbandingan eksperimental dengan baseline diakui sebagai celah penelitian yang penting dan direncanakan sebagai agenda riset lanjutan.

Karena Word2Vec menghasilkan representasi pada tingkat kata, seluruh vektor kata penyusun satu sinopsis dirata-ratakan (Average Word Embeddings) untuk membentuk satu Vektor Dokumen tunggal, baik sebagai Movie Vector (4.000 dokumen data latih) maupun Query Vector (representasi kata kunci pengguna).

2.5 Perhitungan Kemiripan (Cosine Similarity)

Kemiripan semantik antara Query Vector dan tiap Movie Vector pada 1.000 data uji dihitung menggunakan Cosine similarity, yang mengukur orientasi arah antarvektor tanpa dipengaruhi oleh panjang dokumen, sesuai persamaan (1). Film kemudian diurutkan berdasarkan nilai kemiripan tertinggi untuk menghasilkan daftar rekomendasi Top-N. Dengan vA sebagai vektor kata kunci pengguna dan vB sebagai vektor sinopsis film kandidat:

$$\text{Cosine similarity}(vA, vB) = (vA \cdot vB) / (|vA| \times |vB|) \dots(1)$$

2.6 Evaluasi Kinerja

Ground truth dibentuk melalui penilaian dua pakar independen di bidang sinematografi dan literatur fiksi terhadap 25 kueri uji (masing-masing 3 judul teratas, 75 baris data), dan disaring

menggunakan Inter-Annotator Agreement [12] baris data hanya diambil sebagai ground truth apabila kedua pakar sepakat (relevan/tidak relevan), sedangkan baris yang berbeda penilaian dieliminasi. Kinerja algoritma pada jendela Top-3 Recommendation kemudian diukur menggunakan tiga metrik:

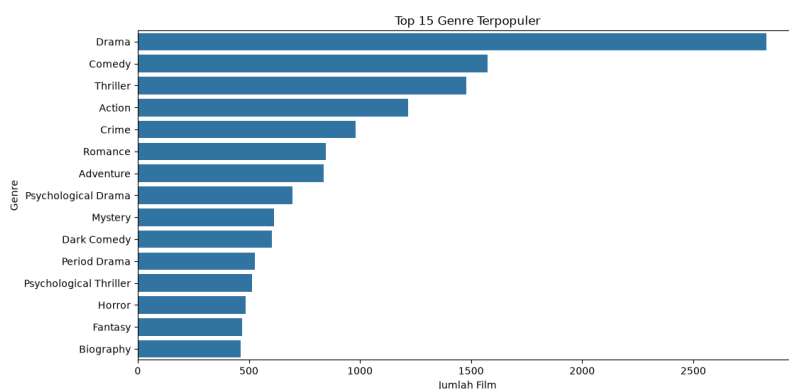
$\text{Precision@K} = \frac{\text{Jumlah film relevan pada K urutan teratas}}{K} \dots(2)$

$\text{Recall@K} = \frac{\text{Jumlah film relevan pada K urutan teratas}}{\text{Total film relevan}} \dots(3)$

$\text{MRR} = \frac{1}{|Q|} \times \sum (1/\text{rank}_i) \dots(4)$

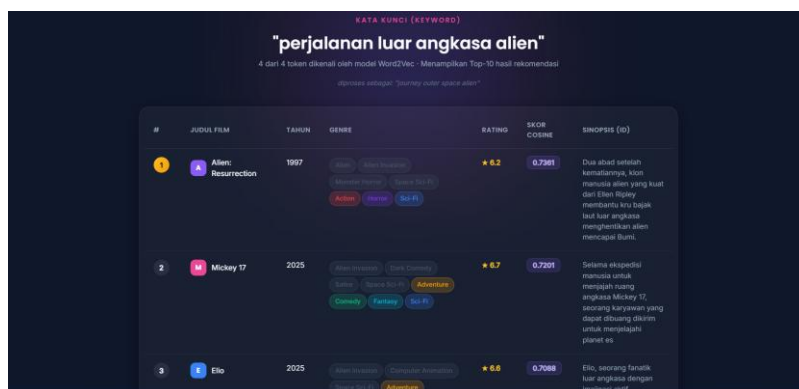
3. HASIL DAN PEMBAHASAN

Proses pengumpulan data menghasilkan populasi akhir sebanyak 5.000 judul film berbahasa Inggris hasil kombinasi dataset publik dan web scraping IMDb, dengan sebaran genre yang beragam sebagaimana ditunjukkan pada Gambar 2.



Gambar 2. Distribusi 15 Kategori Genre Film Terbanyak dalam Dataset

Setelah melalui pembersihan, prapemrosesan, dan pembagian data (4.000 data latih : 1.000 data uji), model Word2Vec Skip-gram dilatih dan menghasilkan representasi vektor berdimensi 200 untuk setiap kata dalam korpus sinopsis. Pengujian Cosine similarity terhadap kata kunci bebas pengguna menghasilkan daftar rekomendasi yang divisualisasikan pada antarmuka prototype web, seperti dicontohkan pada Gambar 3.



Gambar 3. Visualisasi Hasil Rekomendasi pada Antarmuka Aplikasi Web

Validasi ground truth melalui Inter-Annotator Agreement terhadap 75 baris data menghasilkan 67 baris data yang disepakati penuh oleh kedua pakar (18 dari 25 kueri mencapai kesepakatan mutlak di seluruh peringkat), sedangkan 8 baris sisanya dieliminasi karena perbedaan penilaian. 67 baris data bersih inilah yang digunakan sebagai fondasi perhitungan metrik kinerja. Rekapitulasi hasil evaluasi terhadap 25 kueri uji ditunjukkan pada Tabel 2.

Tabel 2. Rekapitulasi Nilai Evaluasi Kinerja Algoritma

Metrik Evaluasi	Nilai
Precision@3	0,5333
Recall@3	0,7800
Mean Reciprocal Rank (MRR)	0,7300

Nilai Precision@3 sebesar 0,5333 (53,33%) menunjukkan bahwa dari 3 film teratas yang disajikan sistem untuk setiap kueri, rata-rata 1 hingga 2 film terbukti relevan dengan maksud pencarian pengguna nilai yang representatif mengingat sistem melakukan pencocokan semantik pada teks sinopsis yang luas, bukan pencarian kata kunci eksak. Nilai Recall@3 yang tinggi (0,7800) membuktikan keunggulan representasi vektor Word2Vec Skip-gram dalam menjangkau sebagian besar film relevan tanpa banyak melewati dokumen penting. Sementara itu, MRR sebesar 0,7300 menunjukkan bahwa film relevan pertama secara rata-rata berhasil ditempatkan pada peringkat pertama atau kedua, yang penting untuk kepuasan pengalaman pencarian pengguna.

Hasil ini menunjukkan bahwa integrasi word embeddings dan Cosine similarity efektif menangkap kemiripan makna naratif meskipun sistem beroperasi murni berdasarkan masukan kata kunci tanpa judul film acuan (seed movie), sejalan dengan temuan bahwa representasi vektor semantik unggul dibandingkan pendekatan berbasis frekuensi kata seperti TF-IDF dalam mengatasi semantic gap [9], [10].

Penelitian ini tidak menyertakan eksperimen perbandingan langsung antara Word2Vec dan TF-IDF maupun model embedding lain (misalnya FastText atau GloVe) dalam satu skenario uji terkontrol. Keterbatasan ini dilatarbelakangi oleh fokus penelitian pada validasi pipeline end-to-end yang mencakup web scraping, pemrosesan teks, dan evaluasi berbasis expert judgment sehingga ruang lingkup dibatasi pada satu konfigurasi algoritma untuk menjaga kedalaman analisis dan kelayakan sumber daya. Perbandingan eksperimental antara Word2Vec dan baseline TF-IDF maupun model embedding alternatif diakui sebagai celah penelitian yang penting dan direncanakan sebagai agenda riset lanjutan.

Terdapat beberapa keterbatasan yang perlu dicatat. Pertama, ground truth dihasilkan dari penilaian dua pakar subjektivitas penilaian relevansi tetap menjadi faktor inheren. Kedua, evaluasi dilakukan pada 25 kueri uji, sehingga generalisasi hasil ke domain kueri yang lebih luas perlu diuji pada skala yang lebih besar. Ketiga, metode Average Word Embeddings tidak mempertimbangkan bobot kepentingan antarkata. Keempat, penelitian ini belum menyertakan evaluasi usabilitas dari perspektif pengguna akhir.

4. KESIMPULAN

Penelitian ini berhasil membangun sistem rekomendasi film berbasis content-based filtering yang mengintegrasikan dataset publik dengan hasil web scraping IMDb, menerapkan model Word2Vec (arsitektur Skip-gram) untuk merepresentasikan makna semantik sinopsis film, dan algoritma Cosine similarity untuk mengukur kedekatan makna terhadap kata kunci bebas pengguna tanpa bergantung pada judul film acuan (seed movie). Evaluasi pada jendela Top-3 Recommendation dengan ground truth tervalidasi pakar menghasilkan Precision@3 sebesar 0,5333, Recall@3 sebesar 0,7800, dan MRR sebesar 0,7300, yang secara keseluruhan menunjukkan kinerja algoritma yang memadai dalam menghasilkan rekomendasi yang relevan secara semantik dan menempatkannya pada peringkat atas, meskipun perbandingan dengan baseline masih diperlukan untuk klaim kinerja yang lebih definitif.

Penelitian selanjutnya disarankan untuk: (1) membandingkan Word2Vec dengan model representasi teks lain seperti FastText, GloVe, atau model berbasis *Transformer* (BERT) guna meningkatkan Precision; (2) mengintegrasikan parameter metadata tambahan (tahun rilis, genre, popularitas) untuk pendekatan *hybrid recommendation*; (3) mengembangkan prototype menjadi

aplikasi web/mobile *production-ready* untuk pengujian usability pengguna akhir dalam skala lebih luas; serta (4) memperluas cakupan evaluasi dengan jumlah kueri dan pakar yang lebih banyak.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada Program Studi Sistem Informasi, Fakultas Teknik, Universitas Muhammadiyah Jember atas dukungan fasilitas dan bimbingan akademik selama penelitian berlangsung. Penulis juga berterima kasih kepada para pakar domain yang telah bersedia memberikan penilaian ground truth secara independen.

PERNYATAAN KONFLIK KEPENTINGAN

Para penulis menyatakan tidak ada konflik kepentingan dalam penelitian ini. Penelitian ini tidak menerima pendanaan dari lembaga publik, komersial, maupun nirlaba mana pun.

DAFTAR PUSTAKA

- [1] S. Badugu and R. Manivannan, "K-Nearest Neighbor and Collaborative Filtering-Based Movie Recommendation System," in *Comput. Networks Inventive Commun. Technol.*, LNDECT, vol. 141, Singapore: Springer, 2023, pp. 461–474, https://doi.org/10.1007/978-981-19-3035-5_35.
- [2] D. P. Kumar et al., "Content Based Recommendation System on Movies," in *Proc. Int. Conf. Emerging Trends Eng. (ICETE 2023)*, 2023, pp. 462–472, https://doi.org/10.2991/978-94-6463-252-1_49.
- [3] Z. Yao, "Review of Movie Recommender Systems Based on Deep Learning," in *Proc. ICLCC 2023*, 2023, p. 02010, <https://doi.org/10.1051/shsconf/202315902010>.
- [4] L. Romero Meza and G. D'Urso, "Filter bubble and Recommendation," *Psychol. Stud. (Mysore)*, vol. 69, no. 3, pp. 349–367, 2024, <https://doi.org/10.1007/s12646-024-00807-0>.
- [5] D. Velamentosa and E. Zuliarso, "Sistem Rekomendasi Film Menggunakan Metode Content-based filtering," *JATI*, vol. 9, no. 2, pp. 2918–2922, 2025, <https://doi.org/10.36040/jati.v9i2.13251>.
- [6] S. Sinha and T. Sharma, "Content-Based Movie Recommendation System: An Enhanced Approach," *Int. J. Innov. Res. Comput. Sci. Technol.*, vol. 11, no. 3, pp. 67–71, 2023, <https://doi.org/10.55524/ijirest.2023.11.3.12>.
- [7] H. Yuan and A. A. Hernandez, "User Cold Start Problem in Recommendation Systems: A Systematic Review," *IEEE Access*, vol. 11, pp. 136958–136977, 2023, <https://doi.org/10.1109/ACCESS.2023.3338705>.
- [8] I. Valova, T. Mladenova, G. Kanev, and T. Halacheva, "Web Scraping - State of Art, Techniques and Approaches," in *Proc. Natl. Conf. Int. Participation TELECOM*, 2023, pp. 1–4, <https://doi.org/10.1109/TELECOM59629.2023.10409723>.
- [9] C. Mediani et al., "Content-Based Recommender System using Word Embeddings for Pedagogical Resources," in *Proc. Int. Conf. Pattern Anal. Intell. Syst. (PAIS)*, 2023, pp. 1–8, <https://doi.org/10.1109/PAIS60821.2023.10321989>.
- [10] S. J. Johnson, M. R. Murty, and I. Navakanth, "A detailed review on word embedding techniques with emphasis on word2vec," *Multimed. Tools Appl.*, vol. 83, no. 13, pp. 37979–38007, 2024, <https://doi.org/10.1007/s11042-023-17007-z>.
- [11] X. Tian, "Content-based Filtering for Improving Movie Recommender System," in *Proc. DAI 2023*, Atlantis Press, Feb. 2024.
- [12] F. Yang, H. Dong, W. Wang, W. Ding, and T. Lin, "A Comprehensive Survey on Inter-Annotator Agreement and the Use of Cohen's Kappa and Fleiss' Kappa for Ground Truth

-
- Validation in AI Algorithm Evaluation," IEEE Access, vol. 11, pp. 21300–21312, 2023, <https://doi.org/10.1109/ACCESS.2023.3249759>.
- [13] A. Kumar and R. Singh, "BERT-Based Semantic Movie Recommendation Using Synopsis Embeddings," J. Inf. Sci. Eng., vol. 40, no. 2, pp. 341–358, 2024, [https://doi.org/10.6688/JISE.202403_40\(2\).0007](https://doi.org/10.6688/JISE.202403_40(2).0007).
- [14] F. Zhang, Y. Liu, and H. Chen, "Hybrid Movie Recommendation Integrating Metadata and Contextual Embeddings," IEEE Access, vol. 12, pp. 15234–15248, 2024, <https://doi.org/10.1109/ACCESS.2024.3358921>.
- [15] I. Prasetya and D. Santoso, "Indonesian-Language Keyword-Based Content Retrieval on Multilingual Film Datasets," J. RESTI, vol. 8, no. 1, pp. 112–121, 2024, <https://doi.org/10.29207/resti.v8i1.5217>.