

**ORIGINAL ARTICLE****A Hybrid Intelligent Recommender System Based on Data Mining for Personalizing Services in Digital Libraries**Peyman Almasinejad <sup>1\*</sup>, Mohammad Zahaby <sup>2</sup>

1. Assistant Professor,  
Department of Knowledge and  
Information Science, Payame  
Noor University, Tehran, Iran.  
2. Assistant Professor,  
Department of Knowledge and  
Information Science, Payame  
Noor University, Tehran, Iran.

**\*Correspondence**

Peyman Almasinejad  
E-mail:  
[p.almasinejad@pnu.ac.ir](mailto:p.almasinejad@pnu.ac.ir)

Receive Date: 25/Mar/2026  
Accept Date: 02/June/2026  
First Publish Date: 02/June/ 2026

**How to cite**

Almasinejad, P., Zahaby, M.  
(2026). A Hybrid Intelligent  
Recommender System Based on  
Data Mining for Personalizing  
Services in Digital Libraries.  
*Digital and Smart Libraries  
Research*, 12(4), 89-104.

**Introduction**

The rapid growth of digital resources has increased the difficulty of identifying relevant information efficiently in digital libraries. Traditional search and filtering approaches are often unable to address users' diverse and personalized information needs. Intelligent recommender systems have therefore emerged as an important solution for improving information discovery and user experience. However, relying solely on content-based or collaborative filtering may lead to limitations such as the cold-start problem, popularity bias, and insufficient recommendation coverage. Accordingly, this study aims to develop a hybrid intelligent recommender system based on data mining to personalize services in digital libraries. The proposed approach combines users' behavioral information with content-based features to generate more accurate, diverse, and reliable recommendations.

**Methodology**

This study developed a hybrid recommendation framework consisting of three main components: data preprocessing, two complementary recommendation engines, and an intelligent aggregation layer. The first engine uses content-based information to identify resources with characteristics similar to users' interests, whereas the second employs collaborative filtering based on users' behavioral patterns and interactions. The outputs of these two engines are subsequently integrated through an intelligent aggregation mechanism to produce the final recommendations. Behavioral data were obtained from the Book-Crossing dataset, while content-based information was extracted from both the Book-Crossing and CiteULike datasets. The performance of the proposed system was evaluated using precision, recall, diversity, coverage, and a simulated user-satisfaction measure. The evaluation focused on comparing the hybrid approach with single-strategy recommendation methods.

**Findings**

The findings demonstrate that the proposed hybrid recommender system performs better than single-strategy recommendation approaches in terms of recommendation accuracy and coverage. Combining behavioral and content-based information enables the system to provide recommendations that are simultaneously more precise and diverse. The results also indicate that the hybrid framework can alleviate important limitations of conventional recommender systems, particularly the cold-start problem and popularity bias. The integration of

data mining with semantic analysis of textual content contributes to more effective modeling of user preferences and resource characteristics. Overall, the experimental results confirm the effectiveness of combining complementary recommendation strategies for improving personalized information services in digital libraries.

#### **Discussion and Conclusion**

The findings suggest that hybrid recommendation provides an effective framework for addressing the complexity of personalized information retrieval in digital libraries. Integrating collaborative and content-based approaches allows the system to exploit both behavioral patterns and the semantic characteristics of digital resources, thereby improving recommendation quality and coverage. From a practical perspective, the proposed framework can help digital libraries deliver more personalized services, improve users' information-search experience, and facilitate more effective utilization of available resources. Its scalable and adaptable structure also provides potential for application in next-generation digital library environments. In conclusion, combining data mining, collaborative filtering, and semantic content analysis can provide a practical foundation for developing intelligent and personalized recommendation services in digital libraries.

#### **KEYWORDS**

Hybrid Recommender System, Data Mining, Digital Libraries, Service Personalization, Collaborative Filtering.



Copyright © 2026, by the author(s). Published by Payame Noor University, Tehran, Iran.

This is an open access article under the CC BY 4.0 license (<http://creativecommons.org/licenses/by/4.0/>).

<https://lib.journals.pnu.ac.ir>

## پژوهش‌های کتابخانه‌های دیجیتال و هوشمند

سال دوازدهم، شماره ۴، پیاپی ۴۷، زمستان ۱۴۰۴ (۸۹-۱۰۴)

DOI: 10.30473/mrs.2026.77677.1698

P-ISSN: 2383-1049

E-ISSN: 2538-5356

«مقاله پژوهشی»

## ارائه یک سیستم توصیه‌گر ترکیبی هوشمند مبتنی بر داده‌کاوی برای شخصی‌سازی خدمات در کتابخانه‌های دیجیتال

پیمان الماسی نژاد<sup>۱\*</sup>، محمد ذهابی<sup>۲</sup>

۱. استادیار، گروه مهندسی کامپیوتر و فناوری اطلاعات، دانشکده فنی و مهندسی، دانشگاه پیام نور، تهران، ایران.  
 ۲. استادیار، گروه مهندسی کامپیوتر و فناوری اطلاعات، دانشکده فنی و مهندسی، دانشگاه پیام نور، تهران، ایران.

\*نویسنده مسئول: پیمان الماسی نژاد

p.almasinejad@pnu.ac.ir

تاریخ دریافت: ۱۴۰۵/۰۱/۰۵

تاریخ پذیرش: ۱۴۰۵/۰۳/۱۲

تاریخ اولین انتشار: ۱۴۰۵/۰۳/۱۲

## استناد به این مقاله:

الماسی نژاد، پیمان و ذهابی، محمد (۱۴۰۵). ارائه یک سیستم توصیه‌گر ترکیبی هوشمند مبتنی بر داده‌کاوی برای شخصی‌سازی خدمات در کتابخانه‌های دیجیتال پژوهش‌های کتابخانه‌های دیجیتال و هوشمند، ۱۲(۴)، ۸۹-۱۰۴.

## چکیده

با رشد فزاینده منابع دیجیتال، کاربران کتابخانه‌های دیجیتال با چالش‌های قابل توجهی در یافتن مؤثر و سریع منابع مرتبط مواجه هستند. روش‌های سنتی جستجو و فیلتر، اغلب نیازهای شخصی‌سازی شده کاربران را برآورده نمی‌کنند و به همین دلیل سیستم‌های توصیه‌گر هوشمند اهمیت ویژه‌ای یافته‌اند. این پژوهش یک سیستم توصیه‌گر ترکیبی ارائه می‌دهد که با بهره‌گیری هم‌زمان از داده‌های رفتاری و مبتنی بر محتوا، خدمات شخصی‌سازی شده را در کتابخانه‌های دیجیتال ارتقا می‌دهد. چارچوب پیشنهادی شامل پیش‌پردازش داده‌ها، دو موتور توصیه‌گر (محتوایی و مشارکتی) و یک لایه ادغام هوشمند است که خروجی‌ها را ترکیب کرده و پیشنهادهای دقیق، متنوع و قابل اعتماد تولید می‌کند. داده‌های رفتاری از مجموعه داده Book-Crossing و اطلاعات محتوایی از منابع Book-Crossing و CiteULike استخراج شد تا ترجیحات کاربران و ویژگی‌های محتوایی منابع مدل‌سازی شود. عملکرد سیستم با معیارهای دقت، پوشش، تنوع و رضایت کاربر مبتنی بر ارزیابی هوشمند ارزیابی شد و نتایج نشان داد که روش ترکیبی نسبت به مدل‌های تک‌رویکرد، دقت و پوشش بالاتری دارد. یافته‌ها نشان می‌دهند که این رویکرد علاوه بر کاهش مشکلاتی مانند «شروع سرد» و سوگیری محبوبیت، موجب بهبود تجربه کاربری و بهینه‌سازی بهره‌برداری از منابع کتابخانه‌های دیجیتال می‌شود و پتانسیل بالایی ترکیب داده‌کاوی و تحلیل معنایی متن را در نسل بعدی بازیابی اطلاعات نشان می‌دهد.

## واژه‌های کلیدی

سیستم توصیه‌گر ترکیبی، داده‌کاوی، کتابخانه‌های دیجیتال، شخصی‌سازی خدمات، فیلترینگ مشارکتی.

حق انتشار این مستند، متعلق به نویسندگان آن است. © ۱۴۰۴ شر این مقاله، دانشگاه پیام نور است.

این مقاله تحت گواهی زیر منتشر شده و هر نوع استفاده غیرتجاری از آن مشروط بر استناد صحیح به مقاله و با رعایت شرایط مندرج در آدرس زیر مجاز است.

This is an open access article under the CC BY (<http://creativecommons.org/licenses/by/4.0/>).<https://lib.journals.pnu.ac.ir>

## مقدمه

با گسترش روزافزون منابع دیجیتال و افزایش حجم اطلاعات در کتابخانه‌های دیجیتال، کاربران با چالشی اساسی در یافتن منابع مرتبط با نیازهای خود مواجه هستند. انبوه مقالات، کتاب‌ها و داده‌های متنوع، در کنار محدودیت توانایی کاربران برای پیمایش مؤثر در این فضای بزرگ، نشان می‌دهد که سیستم‌های سنتی جستجو و فیلتر ساده قادر به تأمین نیازهای دقیق کاربران و شخصی‌سازی تجربه آن‌ها نیستند. (باودن و رایبسون، ۲۰۰۹؛ یاناک و همکاران، ۲۰۱۶). به همین دلیل، اهمیت ارائه راهکارهای هوشمند در این حوزه روبه‌روز بیشتر می‌شود تا کتابخانه‌های دیجیتال بتوانند خدمات سریع و با کیفیت به کاربران ارائه دهند.

سیستم‌های توصیه‌گر به‌عنوان یک راهکار مؤثر معرفی شده‌اند که می‌توانند بر اساس نمایه کاربر و ویژگی‌های منابع، پیشنهادهای متناسب ارائه دهند و رابطی هوشمند بین کاربران و منابع ایجاد کنند (برک، ۲۰۲۲). با این حال، بسیاری از مدل‌های فعلی در کتابخانه‌های دیجیتال تنها از یک رویکرد استفاده می‌کنند: یا فیلتر مبتنی بر محتوا که ویژگی‌های منابع را بررسی می‌کند، یا فیلتر مشارکتی که رفتار کاربران مشابه را ملاک قرار می‌دهد. این مدل‌های تک‌بعدی در مواجهه با چالش‌هایی مانند «شروع سرد»<sup>۴</sup> و پراکندگی علایق کاربران عملکرد ضعیفی دارند (رادپیش، ۲۰۲۱).

در پژوهش‌های اخیر، استفاده از رویکردهای ترکیبی<sup>۵</sup> به‌عنوان یک راه‌حل مؤثر برای غلبه بر محدودیت‌های روش‌های منفرد مطرح شده است (دا سیلوا و همکاران، ۲۰۲۳). این روش‌ها با ادغام اطلاعات محتوایی و رفتاری کاربران، دقت توصیه‌ها را افزایش می‌دهند و تنوع، پوشش و رضایت کاربران را بهبود می‌بخشند (رهنا و همکاران، ۱۳۹۳؛ دیلتون، ۲۰۱۲). با این حال، طراحی چارچوبی مناسب برای کتابخانه‌های دیجیتال که همگام با داده‌کاوی و کاربرپذیری باشد، هنوز نیازمند توسعه است (اسمیت و همکاران، ۲۰۲۰). در این میان، بهره‌گیری از تکنیک‌های داده‌کاوی می‌تواند با استخراج الگوهای رفتاری کاربران و ویژگی‌های محتوایی

منابع، زمینه ارائه توصیه‌های دقیق‌تر و هوشمندتر را در کتابخانه‌های دیجیتال فراهم کند (انصاری و وکیلی‌مفرد، ۲۰۲۱).

مدل پیشنهادی در این مقاله از هر دو منبع مهم داده «رفتار کاربران و محتوای منابع» بهره می‌برد. چارچوب پیشنهادی شامل سه بخش اصلی است: پیش‌پردازش داده‌ها، دو موتور توصیه‌گر مبتنی بر محتوا و مشارکتی، و لایه ترکیب هوشمند برای ادغام نتایج. این ساختار امکان کاهش نقاط ضعف هر روش و ارائه یک سیستم توصیه‌گر دقیق و قابل اعتماد را فراهم می‌کند.

هدف این مطالعه ارائه یک سیستم توصیه‌گر ترکیبی هوشمند مبتنی بر داده‌کاوی برای بهبود ارائه خدمات متناسب با نیازهای اطلاعاتی کاربران در کتابخانه‌های دیجیتال است. ارزیابی سیستم پیشنهادی با معیارهایی مانند دقت، تنوع، پوشش و امتیاز کیفیت توصیه انجام می‌شود تا نشان دهد چگونه ترکیب هوشمند دو رویکرد می‌تواند تجربه کاربری را بهبود بخشد و کارایی کتابخانه‌های دیجیتال را افزایش دهد (جومسری و همکاران، ۲۰۲۴؛ چانو و مورسیو، ۲۰۱۷؛ رهنوی و همکاران، ۲۰۲۲).

در ادامه، مبانی نظری و پیشینه پژوهش مرور شده و سپس روش‌شناسی مدل پیشنهادی تشریح می‌شود. در پایان نیز نتایج ارزیابی و جمع‌بندی پژوهش ارائه می‌گردد.

## مبانی نظری

سیستم‌های توصیه‌گر<sup>۱۰</sup> ابزارهایی هستند که با تحلیل داده‌های کاربران و ویژگی‌های منابع، پیشنهادهای متناسب ارائه می‌دهند و کاربران را در انتخاب منابع مناسب یاری می‌کنند (ریچی و همکاران، ۲۰۱۵). این سیستم‌ها در محیط‌های دیجیتال مانند کتابخانه‌های دیجیتال، فروشگاه‌های آنلاین و پلتفرم‌های رسانه‌ای کاربرد گسترده‌ای دارند.

به‌طور کلی، سه نوع اصلی سیستم توصیه‌گر وجود دارد: فیلتر مبتنی بر محتوا<sup>۱۱</sup>، فیلتر مشارکتی<sup>۱۲</sup>، و سیستم‌های ترکیبی. فیلتر مبتنی بر محتوا پیشنهادها را بر اساس ویژگی‌ها و

10. Ansari & Vakilmofrad

11. Jomsri et al.

12. Çano & Morisio

13. Rhanoui et al.

14. Recommender Systems

15. Ricci et al.

16. Content-Based Filtering

17. Collaborative Filtering

1. Bawden & Robinson

2. Jannach et al.

3. Burke

4. Cold Start

5. Rodpysh

6. Hybrid

7. da Silva et al.

8. De Leon

9. Smith et al.

### پیشینه پژوهش

در سال های اخیر، پژوهش های متعددی در زمینه سیستم های توصیه گر در کتابخانه های دیجیتال انجام شده است. مطالعه ای توسط جومسری و همکاران (۲۰۲۴) مدل سیستم توصیه گر ترکیبی را ارائه داد که ترکیبی از فیلتر مشارکتی و فیلتر مبتنی بر محتوا با وزن دهی ۸۰ - ۲۰ بود. نتایج نشان داد که این ترکیب می تواند دقت و کارایی توصیه ها را نسبت به روش های سنتی افزایش دهد.

همچنین، صبری (۲۰۲۵) در یک مرور جامع به بررسی چالش ها و فرصت های طراحی و پیاده سازی سیستم های توصیه گر ترکیبی در زمینه داده های کلان پرداخت. این مطالعه بر اهمیت توجه به ویژگی های منحصربه فرد داده های کلان مانند حجم، سرعت و تنوع در طراحی سیستم های توصیه گر تأکید دارد.

علاوه بر این، مطالعه ای توسط ابراهیم<sup>۹</sup> (۲۰۲۵) به بررسی روش های مختلف توصیه گر از جمله فیلتر مشارکتی، فیلتر مبتنی بر محتوا و سیستم های ترکیبی پرداخت و چالش هایی مانند مشکل شروع سرد و حباب فیلترینگ را مورد بحث قرار داد.

در حوزه به کارگیری داده کاوی و متن کاوی در سیستم های توصیه گر، وانگ و همکاران<sup>۱۰</sup> (۲۰۲۴) یک سیستم توصیه گر کتاب در کتابخانه های دیجیتال طراحی کردند. این سیستم با تحلیل تاریخچه امانت گیری کاربران و ویژگی های محتوایی کتاب ها، قادر به ارائه پیشنهادهای متناسب با نیاز کاربران است. شو و همکاران<sup>۱۱</sup> (۲۰۲۲) با بهره گیری از تکنیک های داده کاوی در مقیاس وسیع، به تحلیل رفتار کاربران در کتابخانه های دیجیتال پرداختند. هدف اصلی این پژوهش، شناسایی الگوهای پنهان در تعاملات کاربران و بهبود دقت سیستم های توصیه گر بود. نتایج نشان داد که این رویکرد می تواند به ارتقای کیفیت خدمات و بهبود ارائه منابع در محیط های دیجیتال منجر شود.

مطالعه فیلیپ و همکاران<sup>۱۲</sup> (۲۰۲۵) یک روش مبتنی بر محتوا برای توصیه مقالات در کتابخانه های دیجیتال پیشنهاد داد. پژوهش با استخراج ویژگی های معنایی از متن مقالات - با تکنیک هایی مانند TF-IDF<sup>۱۳</sup> و مدل های برداری معنایی -

فیلتر مبتنی بر محتوا پیشنهادها را بر اساس ویژگی ها و محتوای منابع ارائه می دهد و نیازمند تحلیل دقیق متادیتا و متن منابع است (لوپس و همکاران، ۲۰۱۱). فیلتر مشارکتی پیشنهادها را بر اساس رفتار و علایق کاربران مشابه شکل می دهد و از اطلاعات جمعی کاربران بهره می برد (سو و خوش-گفتار، ۲۰۰۹). سیستم های ترکیبی با ادغام دو رویکرد فوق، دقت و تنوع توصیه ها را افزایش می دهند و محدودیت های هر روش منفرد را کاهش می دهند (صبری، ۲۰۲۵).

داده کاوی<sup>۵</sup> و متن کاوی<sup>۶</sup> نقش اساسی در بهبود عملکرد سیستم های توصیه گر ایفا می کنند. داده کاوی با استخراج الگوهای پنهان از رفتار کاربران و تعاملات آن ها، امکان شناسایی علاقه مندی ها و پیش بینی نیازهای آینده را فراهم می کند (هان و همکاران، ۲۰۱۱). متن کاوی، با استخراج ویژگی ها و مفاهیم کلیدی از منابع دیجیتال، امکان ایجاد نمایه دقیق محتوا و افزایش دقت فیلتر مبتنی بر محتوا را فراهم می سازد (مینینگ و همکاران، ۲۰۰۸). ترکیب این دو رویکرد در سیستم های توصیه گر منجر به ارائه پیشنهادهای هوشمندتر و متنوع تر می شود.

سیستم های توصیه گر ابزارهایی هستند که با تحلیل داده های کاربران و ویژگی های منابع، پیشنهادهای متناسب با نیاز اطلاعاتی کاربران ارائه می دهند و آن ها را در انتخاب منابع مناسب یاری می کنند (ریچی و همکاران، ۲۰۱۵). این سیستم ها در محیط های دیجیتال از جمله کتابخانه های دیجیتال، فروشگاه های آنلاین و پلتفرم های رسانه ای کاربرد گسترده ای دارند و نقش مهمی در بهبود دسترسی به اطلاعات ایفا می کنند.

در کتابخانه های دیجیتال، با توجه به حجم بالای منابع و تنوع اطلاعات، استفاده از سیستم های توصیه گر مبتنی بر داده کاوی و متن کاوی می تواند به بهبود دسترسی کاربران به منابع مرتبط، افزایش دقت بازیابی اطلاعات و کاهش زمان جستجو کمک کند. این کاربردها شامل پیشنهاد منابع مشابه، گروه بندی منابع بر اساس موضوع و تحلیل الگوهای رفتاری کاربران برای بهبود خدمات اطلاعاتی کتابخانه است.

1. Collaborative Filtering
2. Lops et al.
3. Su & Khoshgoftaar
4. Sabiri
5. Data Mining
6. Text Mining
7. Han et al.
8. Manning et al.

9. Ibrahim

10. Wang et al.

11. Xu et al.

12. Philip et al.

13. Term Frequency-Inverse Document Frequency

این مطالعه تأکید می‌کند که طراحی مناسب رابط و ارزیابی واقعی کاربران برای دقت توصیه‌ها حیاتی است.

مرور جامع ArXiv (مخزن آنلاین مقالات علمی پیش‌انتشار در حوزه‌های مختلف از جمله هوش مصنوعی و سیستم‌های توصیه‌گر) در سال ۲۰۲۴ پیشرفت‌های نظری و کاربردی در سیستم‌های توصیه‌گر را از سال‌های ۲۰۱۷ تا ۲۰۲۴ بررسی کرد. مطالعه شامل روش‌های مبتنی بر فیلتر مشارکتی، محتوا، گراف و یادگیری عمیق بود و همچنین گرایش‌های نوین مانند استفاده از مدل‌های زبان بزرگ (LLM) را برای کتابخانه‌های دیجیتال تحلیل نمود. نتایج خلأهای پژوهشی و فرصت‌های آینده در این حوزه را نشان داد. این مرور نشان داد که روش‌های ترکیبی و مبتنی بر یادگیری عمیق در سال‌های اخیر بیشترین رشد را در بهبود دقت و شخصی‌سازی توصیه‌ها داشته‌اند، اما همچنان چالش‌هایی مانند مقیاس‌پذیری، شفافیت مدل‌ها و مدیریت داده‌های کم‌تراکم پابرجاست. همچنین نتایج مطالعه تأکید کرد که استفاده از مدل‌های زبان بزرگ و روش‌های گراف‌محور می‌تواند در آینده به افزایش کیفیت توصیه‌ها، درک بهتر زمینه معنایی منابع و ارتقای تعامل کاربران در محیط‌های کتابخانه دیجیتال منجر شود.

مرور پژوهش‌های اخیر نشان می‌دهد که سیستم‌های توصیه‌گر در کتابخانه‌های دیجیتال با ترکیب روش‌های مبتنی بر محتوا و مشارکتی، دقت و پوشش پیشنهادها را بهبود بخشیده‌اند (رهنمویی، ۲۰۲۲؛ جومسری، ۲۰۲۴). بهره‌گیری از داده‌کاوی و متن‌کاوی نیز امکان شناسایی الگوهای پنهان و بهبود ارائه خدمات را فراهم کرده است (شو، ۲۰۲۲؛ فیلیپ، ۲۰۲۵). با این حال، مدل‌های فعلی با چالش‌هایی مانند مشکل شروع سرد، سوگیری موقعیتی، محدودیت پوشش و مقیاس‌پذیری محدود روبه‌رو هستند (بیل، ۲۰۱۸؛ یاداو، ۲۰۲۲). این شرایط نشان می‌دهد که نیاز به چارچوبی یکپارچه وجود دارد که دقت، تنوع و تجربه کاربر را همزمان ارتقا دهد (صبری، ۲۰۲۵ و ArXiv, 2024).

بررسی پژوهش‌های پیشین نشان می‌دهد که سیستم‌های توصیه‌گر در کتابخانه‌های دیجیتال با ترکیب روش‌های مبتنی بر محتوا و مشارکتی، دقت و پوشش پیشنهادها را افزایش داده‌اند. همچنین بهره‌گیری از داده‌کاوی و متن‌کاوی امکان شناسایی الگوهای پنهان و بهبود ارائه خدمات را فراهم کرده است. با این حال، چالش‌هایی مانند مشکل شروع سرد،

هونگ و همکاران<sup>۱</sup> (۲۰۰۲) یک سیستم توصیه‌گر مبتنی بر گراف برای کتابخانه‌های دیجیتال معرفی کردند. در این روش، روابط میان کاربران و منابع به صورت گراف مدل‌سازی شد و از شبکه‌های عصبی هاپفیلد<sup>۲</sup> برای تحلیل این روابط بهره گرفته شد. نتایج نشان داد که رویکرد گرافی می‌تواند تنوع و دقت توصیه‌ها را به شکل معناداری بهبود دهد.

رهنویی و همکاران<sup>۳</sup> (۲۰۲۲) سامانه‌ای توصیه‌گر ترکیبی برای پشتیبانی از تصمیم‌گیری خرید و حذف منابع<sup>۴</sup> در کتابخانه‌ها ارائه کردند. این مدل با ترکیب یادگیری ماشین و تحلیل نظرات و امتیازهای کاربران، الگوهای استفاده را استخراج کرده و پیشنهادهای خرید - حذف را تولید می‌کند. ارزیابی موردی در یک کتابخانه ملی بهبود کارایی توسعه مجموعه را نشان داد.

کرویتس و همکاران<sup>۵</sup> (۲۰۲۲) در یک مرور نظام‌مند به بررسی سامانه‌های توصیه‌گر مقالات علمی پرداختند. آن‌ها رویکردهای متداول شامل فیلتر مشارکتی، مبتنی بر محتوا و روش‌های ترکیبی را تحلیل کرده و معیارهای ارزیابی همچون دقت، پوشش و تنوع را مرور نمودند. این مطالعه همچنین بر چالش‌هایی مانند شروع سرد و سوگیری داده‌ها در محیط‌های کتابخانه دیجیتال تأکید کرد.

یاداو و همکاران<sup>۶</sup> (۲۰۲۲) الگوریتم SPACE-R<sup>۷</sup> را برای سیستم‌های توصیه‌گر آکادمیک و کتابخانه‌ای معرفی کردند. این روش با هدف بهبود مقیاس‌پذیری و دقت در مجموعه‌های بزرگ، تعامل کاربران و ویژگی‌های منابع را تحلیل می‌کند. نتایج نشان داد الگوریتم پیشنهادی کارایی بالاتر و پاسخ‌دهی سریع‌تر نسبت به روش‌های پایه دارد و برای کتابخانه‌های دیجیتال مناسب است.

بیل و همکاران<sup>۸</sup> (۲۰۱۸) سوگیری موقعیت<sup>۹</sup> را بر تعامل کاربران در سیستم‌های توصیه‌گر علمی بررسی کردند. نتایج نشان داد که ترتیب نمایش آیت‌ها می‌تواند سوگیری محبوبیت ایجاد کرده و ارزیابی ساده سیستم‌ها را تحت تأثیر قرار دهد.

1. Huang et al.
2. Hopfield Network
3. Rhanoui et al.
4. Acquisition and Weeding
5. Kreutz et al.
6. Yadav et al.
7. Scalable Personalized Academic/Content Engine – Recommendation
8. Beel et al.
9. Position Bias

دیجیتال نیازمند رویکردی است که علاوه بر دقت، بتواند تنوع و پوشش اطلاعاتی را نیز به صورت متوازن تأمین کند. تمرکز این تحقیق بر ارائه توصیه های دقیق و قابل اعتماد است تا کاربران بتوانند منابع مورد نیاز خود را به شکل سریع و مؤثر بیابند.

چارچوب کلی مدل شامل چهار بخش اصلی است: ورودی ها که شامل داده های کاربران و ویژگی های منابع می باشد، مرحله پیش پردازش برای پاک سازی داده ها و استخراج ویژگی ها، دو موتور توصیه گر مجزا (مبتنی بر محتوا و مشارکتی) برای تولید پیشنهادهای، و لایه ادغام و رتبه بندی هوشمند که نتایج هر موتور را ترکیب کرده و پیشنهادهای نهایی را ارائه می دهد. این ساختار موجب کاهش ضعف های هر روش منفرد، افزایش پوشش و تنوع توصیه ها و ایجاد پایه ای مناسب برای تحلیل های بعدی می شود.

### ورودی ها و داده ها

روش پیشنهادی این پژوهش بر پایه ی ترکیب داده های رفتاری و محتوایی طراحی شده است؛ از این رو، نیازمند دو منبع داده ای مجزا اما مکمل است. نخستین منبع به تحلیل رفتار کاربران در محیط های کتابخانه ای می پردازد و دومین، ارتباط بین مقالات علمی را پوشش می دهد.

از آنجا که مدل پیشنهادی این پژوهش به طور همزمان به توصیه کتاب ها و مقالات علمی می پردازد، لازم است دو نوع داده مستقل اما مکمل در اختیار داشته باشیم؛ یکی برای تحلیل تعاملات کاربر با منابع کتابی و دیگری برای تحلیل محتوای متون علمی. این تفکیک داده ای باعث می شود مدل بتواند دامنه وسیع تری از نیازهای اطلاعاتی کاربران کتابخانه های دیجیتال را پوشش دهد و دقیق تر ارائه کند. از آنجا که مدل پیشنهادی این پژوهش به طور همزمان به توصیه کتاب ها و مقالات علمی می پردازد، لازم است دو نوع داده مستقل اما مکمل در اختیار داشته باشیم؛ یکی برای تحلیل تعاملات کاربر با منابع کتابی و دیگری برای تحلیل محتوای متون علمی. این تفکیک داده ای باعث می شود مدل بتواند دامنه وسیع تری از نیازهای اطلاعاتی کاربران کتابخانه های دیجیتال را پوشش دهد و پیشنهادهایی جامع تر و دقیق تر ارائه کند.

در بخش کتاب ها، از مجموعه داده Book-Crossing (BX) استفاده شده که شامل سه فایل مکمل است: BX-Book-Ratings داده های رفتاری کاربران را در قالب امتیازدهی به کتاب ها نگهداری می کند، BX-Books حاوی اطلاعات متنی و محتوایی کتاب ها (عنوان، نویسنده، ناشر و

سوگیری موقعیتی، محدودیت پوشش و مقیاس پذیری همچنان وجود دارد. در پاسخ به این محدودیت ها، مقاله حاضر چارچوبی ترکیبی و هوشمند ارائه می دهد که با ادغام تحلیل محتوایی و داده های رفتاری، دقت، تنوع و تجربه کاربر را ارتقا می بخشد و خلأهای پژوهش های قبلی را پوشش می دهد.

### روش شناسی پژوهش و مدل پیشنهادی

پژوهش حاضر با هدف ارائه یک سیستم توصیه گر ترکیبی، در قالب یک مطالعه کاربردی انجام شده است. این تحقیق تلاش دارد دقت، تنوع و تجربه کاربری در کتابخانه های دیجیتال را بهبود بخشد. در این راستا، با بهره گیری از داده های تاریخی کاربران و ویژگی های محتوایی منابع، چارچوبی ترکیبی شامل پیش پردازش داده ها، دو موتور توصیه گر (مبتنی بر محتوا و مشارکتی) و یک لایه ادغام هوشمند ارائه شده است. هدف اصلی پژوهش، طراحی مدلی است که خلأهای موجود در تحقیقات پیشین را پوشش داده و خدمات متناسب با نیاز کاربران را به صورت قابل اتکا ارائه کند. مدل پیشنهادی این پژوهش به صورت چنددانه<sup>۱</sup> طراحی شده و قادر است به طور همزمان کتاب ها و مقالات علمی را در قالب یک چارچوب یکپارچه توصیه کند. جزئیات مدل و الگوریتم ها در بخش های بعدی تشریح شده اند.

### هدف و چارچوب کلی مدل

هدف اصلی این پژوهش، طراحی و ارائه مدلی ترکیبی برای سیستم های توصیه گر در کتابخانه های دیجیتال است که بتواند دقت، پوشش و تنوع توصیه ها را بهبود بخشد و تجربه کاربران را ارتقا دهد. مدل پیشنهادی با بهره گیری از داده های تاریخی کاربران و ویژگی های محتوایی منابع، تلاش می کند محدودیت های مدل های تک رویکردی از جمله مسئله شروع سرد و کاهش پوشش را کاهش دهد.

شایان ذکر است که نیازهای اطلاعاتی در کتابخانه های دیجیتال ماهیتی متفاوت از سامانه های تجاری مانند آمازون<sup>۲</sup> (پلتفرم تجارت الکترونیک و سیستم توصیه گر تجاری مبتنی بر رفتار کاربران) دارند؛ در محیط های تجاری، هدف اصلی اغلب افزایش فروش و پیشنهاد محصولات مشابه بر اساس رفتار خرید کاربران است، در حالی که در کتابخانه های دیجیتال تمرکز بر دسترسی علمی، تنوع منابع، پوشش جامع موضوعی و حمایت از اهداف پژوهشی و آموزشی کاربران قرار دارد. از این رو، طراحی سیستم های توصیه گر در کتابخانه های

1. Multi-domain  
2. Amazon

مخازن معتبر علمی تأیید شده است. مجموعه Book-Crossing یکی از شناخته‌شده‌ترین دیتاست‌های حوزه سیستم‌های توصیه‌گر کتاب است که در پژوهش‌های متعددی برای تحلیل رفتار کاربران و ارزیابی الگوریتم‌های توصیه‌گر مورد استفاده قرار گرفته است (زیکر و همکاران، ۲۰۰۴؛ گروور و همکاران، ۲۰۲۲). همچنین مجموعه CiteULike به‌عنوان یک پایگاه داده معتبر در حوزه توصیه مقالات علمی شناخته می‌شود و در مطالعات مرتبط با سیستم‌های توصیه‌گر علمی، تحلیل علاقه‌مندی پژوهشگران و مدل‌سازی تعاملات کاربر - مقاله کاربرد گسترده‌ای داشته است (وانگ و بلی، ۲۰۱۱؛ کرویتس و همکاران، ۲۰۲۲). استفاده از این دیتاست‌ها به دلیل اعتبار پژوهشی، دسترسی پذیری عمومی و کاربرد مکرر در تحقیقات پیشین، قابلیت اعتماد مدل پیشنهادی را افزایش می‌دهد. در جدول ۱، مشخصات اصلی هر مجموعه داده و نقش آن در مدل پیشنهادی نمایش داده شده است.

جدول ۱. مشخصات مجموعه داده‌های مورد استفاده در پژوهش

| نام دیتاست      | نوع داده              | منبع / دسترسی         | محتوای اصلی                                  | نقش در مدل                |
|-----------------|-----------------------|-----------------------|----------------------------------------------|---------------------------|
| BX-Book-Ratings | رفتاری (Behavioral)   | عمومی (Kaggle)        | امتیازدهی کاربران به کتاب‌ها                 | تحلیل ترجیحات کاربران     |
| BX-Books        | محتوایی (Content)     | عمومی (Kaggle)        | اطلاعات متنی کتاب‌ها (عنوان، نویسنده و ناشر) | استخراج ویژگی‌های محتوایی |
| BX-Users        | زمینه‌ای (Contextual) | عمومی (Kaggle)        | اطلاعات کاربران (سن، کشور، محل سکونت)        | غنی‌سازی پروفایل کاربر    |
| CiteULike       | محتوایی و رفتاری      | عمومی (citeulike.org) | داده‌های برچسب‌گذاری و علاقه‌مندی مقالات     | توصیه مقالات مرتبط        |

قرار گرفته و امکان مقایسه‌پذیری بین کاربران فراهم شود. همچنین برای کاهش اثر داده‌های پرت، مقادیر غیرعادی در فرآیند پیش‌پردازش شناسایی و حذف شدند. برای اطمینان از کیفیت داده‌ها، تنها کاربرانی که حداقل سه تعامل معتبر با منابع داشتند حفظ شدند. این آستانه به‌منظور حذف کاربران با تعاملات محدود یا تصادفی و افزایش پایداری پروفایل کاربران انتخاب شده است. در این پژوهش، «تعامل معتبر» به هرگونه امتیازدهی ثبت‌شده و غیرتهی کاربران به منابع در بازه مجاز سیستم اطلاق می‌شود. خروجی این مرحله به صورت ماتریسی از تعاملات کاربر - منبع ذخیره گردید که مبنای یادگیری در بخش مشارکتی مدل است. شکل ۱ روند کلی این فرآیند را نمایش می‌دهد.

کلیدواژه‌ها) است، و BX-Users ویژگی‌های زمینه‌ای کاربران مانند موقعیت جغرافیایی و سن را ذخیره می‌کند. این مجموعه داده از نسخه عمومی Book-Crossing Dataset (در دسترس از طریق Kaggle / UCI Repository) استفاده شده است.

این ساختار سه‌گانه امکان مدل‌سازی دقیق‌تر ترجیحات کاربران و تحلیل متقابل رفتار و محتوا را فراهم می‌آورد.

برای بخش مقالات، از مجموعه‌ی CiteULike بهره گرفته شده است. این دیتاست شامل سوابق واقعی علاقه‌مندی پژوهشگران، کلیدواژه‌ها و ارتباط میان مقالات است و به صورت عمومی در دسترس قرار دارد. انتخاب این دو منبع داده، به دلیل جامعیت، تنوع و تناسب با اهداف پژوهش انجام شده است.

اعتبار مجموعه داده‌های مورد استفاده در این پژوهش از طریق کاربرد گسترده آن‌ها در مطالعات پیشین و انتشار در

### پیش‌پردازش داده‌ها

پیش‌پردازش داده‌ها مرحله‌ای کلیدی در طراحی سیستم‌های توصیه‌گر است که کیفیت و انسجام داده‌ها را برای آموزش مدل تضمین می‌کند. از آنجا که مدل پیشنهادی ما ترکیبی از دو منبع اطلاعاتی متفاوت است - یعنی داده‌های رفتاری کاربران و داده‌های محتوایی منابع - فرآیند پیش‌پردازش نیز باید به‌گونه‌ای طراحی شود که هر دو جنبه‌ی تعامل و معنا را به‌طور هم‌زمان پوشش دهد. این تفکیک داده‌ای، دقت مدل را در تحلیل ترجیحات کاربران و شباهت‌های محتوایی افزایش می‌دهد.

در گام نخست، داده‌های رفتاری از فایل‌های BX-Users و BX-Book-Ratings استخراج و پاک‌سازی شدند.

داده‌های ناقص، تکراری یا نامعتبر حذف و امتیازها به روش نرمال‌سازی حداقل - حداکثر نرمال‌سازی گردیدند تا مقادیر در بازه [۰،۱]



شکل ۱. فرایند پیش پردازش داده‌های رفتاری کاربران در سیستم توصیه گر ترکیبی

ضعیف است. Word2Vec این مشکل را با یادگیری ارتباط معنایی میان کلمات تا حدی رفع می‌کند، ولی هنگام نمایش کل سند، بخشی از ساختار متنی از بین می‌رود (ذهابی و همکاران، ۲۰۲۴). روش Doc2Vec با مدل سازی سطح سند و حفظ روابط معنایی در سطح جمله و پاراگراف، توانایی بالاتری در بازنمایی شباهت‌های واقعی میان متون دارد (لی و میکولوف<sup>۳</sup>، ۲۰۱۴). بنابراین در این پژوهش Doc2Vec به عنوان روش نهایی انتخاب شد، چرا که دقت و کارایی بالاتری در محیط‌های علمی و متنی ارائه می‌دهد. شکل ۲ فرایند کامل استخراج و بردار سازی متون را نشان می‌دهد.

در مرحله بعد، داده‌های محتوایی از توضیحات کتاب‌ها (در مجموعه‌ی BX-Books) و چکیده‌ی مقالات علمی (از پایگاه CiteULike) استخراج شدند. ابتدا متون پاک سازی شدند؛ علائم نگارشی و نویسه‌های زائد حذف، واژه‌ها در متن انگلیسی به حالت پایه بازگردانده شدند<sup>۱</sup> و کلمات توقف<sup>۲</sup> کنار گذاشته شد. این گام موجب حذف نویز و یکنواختی زبانی داده‌ها شد و زمینه را برای استخراج ویژگی‌های عددی فراهم کرد.

در مرحله‌ی استخراج ویژگی، سه روش رایج مورد بررسی قرار گرفت: TF-IDF، Word2Vec و Doc2Vec. روش TF-IDF اگرچه ساده و قابل تفسیر است، اما در درک معنایی واژه‌ها



شکل ۲. فرایند پیش پردازش داده‌های محتوایی در سیستم توصیه گر ترکیبی

1. Lemmatization  
2. Stop Words  
3. Le & Mikolov

آماده شود. این ساختار، امکان یادگیری هم‌زمان از رفتار کاربران و محتوای متون را فراهم می‌کند و مبنای اصلی مدل پیشنهادی در مرحله‌ی آموزش و ارزیابی محسوب می‌شود. بردارهای عددی ثابت دارد و بدین ترتیب، مقایسه دقیق‌تری میان اسناد علمی یا توصیف‌های کتاب‌ها فراهم می‌آورد.

خروجی این بخش به‌جای تولید نمرات شباهت به‌طور مستقیم در ماتریس شباهت، مبتنی بر کسینوس<sup>۱</sup> ذخیره می‌شود. در این ماتریس، بردارهای ویژگی استخراج‌شده از هر منبع (کتاب یا مقاله) در قالب نمایش‌های عددی در نظر گرفته شده و میزان شباهت بین منابع مختلف بر اساس فاصله کسینوسی محاسبه می‌گردد. این ماتریس به‌طور مؤثر شباهت‌های معنایی میان منابع را مدل‌سازی کرده و می‌تواند به‌عنوان ورودی برای مراحل بعدی سیستم توصیه‌گر استفاده شود.

ماتریس شباهت‌ها نقش کلیدی در بهبود دقت توصیه‌ها ایفا می‌کند، زیرا به الگوریتم این امکان را می‌دهد که علاوه بر منابعی که کاربر قبلاً به آن‌ها علاقه‌مند بوده است، منابع جدیدی که ویژگی‌های مشابه‌تری دارند نیز به‌طور مؤثر پیشنهاد کند. در ادامه، این ماتریس به صورت ترکیبی با ماتریس تعاملات رفتاری کاربران و ویژگی‌های محتوایی منابع در قالب یک رویکرد ترکیب خطی وزن‌دار<sup>۲</sup> ادغام می‌شود. در این روش، خروجی هر منبع از موتورهای مبتنی بر محتوا و فیلتر مشارکتی با ضرایب وزنی از پیش تعیین‌شده ترکیب شده و یک امتیاز نهایی برای هر آیتام محاسبه می‌گردد. این فرآیند امکان ادغام هم‌زمان اطلاعات رفتاری و محتوایی را فراهم کرده و منجر به تولید توصیه‌های دقیق‌تر می‌شود. در شکل ۳، ماتریس شباهت‌ها بین ویژگی‌های منابع مختلف نمایش داده شده است که نحوه محاسبه شباهت‌ها و تأثیر آن‌ها بر تولید توصیه‌های نهایی را نشان می‌دهد.

در پایان، خروجی هر دو مرحله - یعنی ماتریس تعاملات کاربران و بردارهای محتوایی منابع - با یک کلید مشترک (شناسه منبع) ادغام شدند تا داده‌ی نهایی برای مدل ترکیبی

### مدل پیشنهادی و الگوریتم توصیه

در این بخش، چارچوب کلی مدل پیشنهادی تشریح می‌شود: مدل یک سامانه توصیه‌گر ترکیبی است که هم‌زمان از اطلاعات محتوایی منابع و الگوهای رفتاری کاربران استفاده می‌کند تا پیشنهادهایی دقیق‌تر، متنوع‌تر و با پوشش وسیع‌تر ارائه دهد. هدف طراحی کاهش ضعف‌های روش‌های تک‌رویکردی - از جمله مشکل «شروع سرد» و سوگیری محبوبیت - و افزایش تجربه کاربر است. معماری کلی شامل سه لایه اصلی است: (۱) پیش‌پردازش و استخراج ویژگی‌ها از داده‌های محتوایی و رفتاری که نورخان‌تر مورد بررسی قرار گرفت، (۲) دو موتور مستقل تولید پیشنهاد (مبتنی بر محتوا و مشارکتی)، و (۳) یک لایه ترکیب هوشمند که نتایج را ادغام و رتبه‌بندی نهایی را تولید می‌کند. در پاراگراف بعدی موتور مبتنی بر محتوا تشریح خواهد شد.

در لایه دوم، موتور مبتنی بر محتوا هدف دارد تا منابع جدید (اعم از کتاب‌ها یا مقالات) بر اساس شباهت محتوایی به منابعی که کاربر پیش‌تر مطالعه کرده، شناسایی و توصیه شوند. برای سنجش این شباهت، از ویژگی‌های استخراج‌شده در مرحله‌ی پیش‌پردازش استفاده می‌شود. در اینجا از الگوریتم Doc2Vec بهره گرفته شده است، زیرا این روش برخلاف TF-IDF که صرفاً بر پایه‌ی فراوانی واژه‌ها عمل می‌کند، می‌تواند معنای ضمنی و بافت مفهومی جملات را نیز در نظر بگیرد. همچنین نسبت به Word2Vec که سطح کلمه را مینا قرار می‌دهد، Doc2Vec قابلیت نمایش کل اسناد را در قالب

جدول ۲. ماتریس شباهت‌ها بین ویژگی‌های استخراج‌شده و منابع در سیستم توصیه‌گر مبتنی بر محتوا

| ویژگی‌ها/منبع | کتاب ۱ | کتاب ۲ | مقاله ۱ | مقاله ۲ |
|---------------|--------|--------|---------|---------|
| ویژگی ۱       | ۰/۸    | ۰/۶    | ۰/۷     | ۰/۵     |
| ویژگی ۲       | ۰/۷    | ۰/۸    | ۰/۶     | ۰/۶     |
| ویژگی ۳       | ۰/۹    | ۰/۷    | ۰/۸     | ۰/۷     |
| ویژگی ۴       | ۰/۶    | ۰/۷    | ۰/۹     | ۰/۸     |

منابع - نظیر امتیازدهی، تاریخچه مطالعه و الگوهای استفاده - مبنای تحلیل قرار می‌گیرد.

در همان لایه دوم، موتور توصیه‌گر مبتنی بر فیلترینگ مشارکتی به کار گرفته می‌شود تا الگوهای رفتاری مشترک میان کاربران شناسایی شود. در این بخش، تعاملات کاربران با

1. Cosine Similarity Matrix  
2. Weighted Linear Combination

فیلترینگ مشارکتی با تحلیل الگوهای تعامل میان کاربران و منابع، قادر به شناسایی شباهت‌های رفتاری و استخراج ترجیحات ضمنی کاربران است. با این حال، این روش در شرایطی که داده‌های تعامل کافی برای کاربران جدید یا منابع تازه وجود نداشته باشد، با محدودیت مواجه می‌شود و عملکرد آن به میزان داده‌های موجود وابسته است. از این رو، در مدل پیشنهادی، این روش در کنار موتور مبتنی بر محتوا استفاده شده تا اثر محدودیت‌های آن کاهش یابد. نمونه‌ای از ماتریس تعامل کاربران و منابع پیشنهادی در جدول ۳ نمایش داده شده است.

جدول ۳. ماتریس تعامل کاربران و منابع پیشنهادی

| کتاب   | میانگین امتیاز |
|--------|----------------|
| کتاب A | ۴/۵            |
| کتاب B | ۴/۲            |
| کتاب C | ۴/۶            |

پس از اعمال وزن‌ها، هر موتور خروجی خود را تولید می‌کند. موتور محتوایی منابع مشابه را بر اساس ویژگی‌های محتوایی منابع شبیه‌سازی می‌کند، در حالی که موتور فیلتر مشارکتی به دنبال شبیه‌سازی رفتار کاربران مشابه است و منابعی که کاربران مشابه به آن‌ها علاقه‌مند بوده‌اند را پیشنهاد می‌دهد. سپس خروجی‌ها به‌طور هوشمند ترکیب می‌شوند. در این مرحله، وزن‌های هر موتور به‌طور ثابت در نظر گرفته شده و نمرات منابع به‌طور متناسب با وزن هر موتور ترکیب می‌شوند. نمره نهایی برای هر منبع به دست می‌آید که ترکیبی از نمرات هر دو موتور است. این نمره نهایی به‌عنوان ورودی به بخش پیشنهادی نهایی سیستم توصیه گر ارائه می‌شود. روش ترکیب مطابق فرمول شماره ۱ انجام می‌شود.

هدف این موتور، استخراج روابط پنهان میان کاربران و منابع از طریق بررسی شباهت‌های رفتاری است. برای این منظور، کاربران مشابه بر اساس میزان شباهت در الگوهای تعامل و امتیازدهی تعیین می‌شوند. شباهت میان کاربران با استفاده از معیار شباهت کسینوسی میان بردارهای تعامل کاربر - منبع محاسبه شده و کاربرانی که بیشترین میزان شباهت را دارند به‌عنوان کاربران مشابه انتخاب می‌شوند. سپس منابعی که توسط این کاربران مشابه امتیاز بالاتری دریافت کرده‌اند، به‌عنوان گزینه‌های پیشنهادی در نظر گرفته می‌شوند.

در لایه سوم، یا همان لایه ترکیب، خروجی‌های حاصل از موتورهای محتوایی و فیلتر مشارکتی به صورت خودکار ترکیب می‌شوند تا توصیه‌های دقیق‌تری به کاربر ارائه گردد. هدف این لایه، بهبود کیفیت توصیه‌ها از طریق ادغام مؤثر خروجی‌های این دو موتور است.

در لایه ترکیب، برای ادغام خروجی موتورهای محتوایی و مشارکتی از یک روش ترکیب خطی وزن‌دار استفاده شده است. در این روش، وزن هر موتور در مرحله طراحی مدل و بر اساس ارزیابی اولیه عملکرد آن تعیین می‌شود. معیارهایی مانند دقت، بازخوانی، پوشش و تنوع در تعیین سهم هر موتور مورد توجه قرار گرفته‌اند. پس از تعیین اولیه، این وزن‌ها در مرحله اجرای سیستم به صورت ثابت اعمال می‌شوند تا فرآیند ترکیب نتایج با ساختاری پایدار و قابل تکرار انجام گیرد.

$$(collaborative^S \times collaborative^W) + (content^S \times content^W) = final^S \quad (1)$$

که در آن:

- $final^S$  نمره نهایی پیشنهادی است.
  - $collaborative^W$  و  $content^W$  وزن‌های موتور محتوایی و مشارکتی هستند.
  - $collaborative^S$  و  $content^S$  نمرات خروجی موتورهای محتوایی و مشارکتی برای یک منبع خاص هستند.
- این فرمول به‌طور خودکار نمره نهایی را از ترکیب خروجی‌های دو موتور محاسبه کرده و بهترین پیشنهاد را به کاربر می‌دهد. این ترکیب به سیستم این امکان را می‌دهد که بدون نیاز به فیدبک

مداوم از کاربر، بهترین منابع را بر اساس ترکیب ویژگی‌های محتوایی و الگوهای رفتاری کاربران مشابه به کاربر پیشنهاد دهد. سیستم به‌طور دائم و خودکار با استفاده از این وزن‌ها، به‌طور هوشمندانه ترکیب موتورهای محتوایی و مشارکتی را انجام می‌دهد و به کاربر منابع مناسب‌تری پیشنهاد می‌کند.

شکل ۳، مدل شامل سه موتور مختلف (محتوایی، فیلتر مشارکتی و ترکیب هوشمند) را به تصویر می‌کشد که از داده‌های تعامل کاربران و منابع برای تولید توصیه‌های نهایی استفاده می‌کند.



شکل ۳. مدل سیستم توصیه‌گر ترکیبی

جنبه‌ای از عملکرد سیستم را ارزیابی کرده و به بهبود کیفیت توصیه‌ها کمک می‌کنند.

دقت<sup>۲</sup>: دقت نشان‌دهنده نسبت اقلام مرتبط در میان اقلام پیشنهادی سیستم است. این شاخص مشخص می‌کند چه میزان از پیشنهادهای ارائه‌شده واقعاً با ترجیحات یا تعاملات واقعی کاربران مطابقت دارند. مقدار بالاتر دقت بیانگر آن است که سیستم توانسته اقلام مرتبط‌تری را در فهرست توصیه‌ها قرار دهد (ذهابی و همکاران، ۲۰۲۵؛ هرلاکر و همکاران، ۲۰۰۴؛ ریچی و همکاران، ۲۰۱۵).

بازخوانی<sup>۳</sup>: بازخوانی نشان‌دهنده توانایی سیستم در بازیابی اقلام مرتبط از میان کل اقلام مرتبط موجود است. این معیار بیان می‌کند چه نسبتی از منابعی که واقعاً برای کاربر مناسب بوده‌اند، توسط سیستم شناسایی و پیشنهاد شده‌اند. مقدار بالاتر بازخوانی نشان‌دهنده کاهش از دست رفتن اقلام مرتبط در فرآیند توصیه است (ذهابی و همکاران، ۲۰۲۵؛ هرلاکر و همکاران، ۲۰۰۴؛ ریچی و همکاران، ۲۰۱۵).

پوشش<sup>۷</sup>: پوشش به میزان توانایی سیستم در ارائه پیشنهاد برای بخش گسترده‌ای از اقلام موجود در مجموعه داده اشاره دارد. این شاخص نشان می‌دهد سیستم تا چه حد می‌تواند منابع متنوعی را از کل فضای آیتم‌ها وارد فرآیند توصیه کند و تنها به

در مدل پیشنهادی، مفاهیم داده‌کاوی در فرآیند استخراج دانش از داده‌های محتوایی و رفتاری به کار گرفته شده‌اند. در بخش محتوایی، از الگوریتم Doc2Vec برای استخراج ویژگی‌های معنایی منابع و تبدیل اسناد به بردارهای عددی استفاده شده است. این نمایش برداری امکان محاسبه شباهت میان منابع مختلف را بر اساس ویژگی‌های معنایی فراهم می‌کند. همچنین در بخش فیلترینگ مشارکتی، تحلیل داده‌های رفتاری کاربران و تعاملات ثبت‌شده میان کاربران و منابع، مبنای شناسایی الگوهای ترجیحی قرار می‌گیرد. در این بخش، شباهت کاربران بر اساس الگوهای امتیازدهی و تعاملات گذشته محاسبه می‌شود تا منابع مرتبط شناسایی شوند. به‌طور کلی، استفاده از روش‌های داده‌کاوی در این پژوهش امکان استخراج روابط پنهان از داده‌ها و بهبود فرآیند پیشنهاددهی را فراهم می‌سازد. در این پژوهش، منظور از داده‌کاوی استفاده از تکنیک‌های بردارسازی<sup>۱</sup> و تحلیل شباهت مبتنی بر کسینوس برای استخراج روابط پنهان میان منابع و کاربران است.

### معیارهای ارزیابی

در ارزیابی عملکرد سیستم‌های توصیه‌گر، انتخاب شاخص‌های مناسب برای سنجش دقت، کیفیت و توانایی سیستم در ارائه پیشنهادهای مؤثر اهمیت ویژه‌ای دارد. این شاخص‌ها امکان ارزیابی دقیق عملکرد سیستم را فراهم کرده و به درک بهتر توانایی آن در مواجهه با داده‌ها و ترجیحات کاربران کمک می‌کنند. از جمله این شاخص‌ها می‌توان به دقت، پوشش، تنوع و امتیاز کیفیت توصیه اشاره کرد. هر یک از این معیارها

2. Precision  
3. Zahaby et al.  
4. Herlocker et al.  
5. Ricci et al.  
6. Recall  
7. Coverage

1. Doc2Vec

ماشین استفاده شده است. در بخش پیش پردازش داده ها و تحلیل ماتریس های رفتاری، کتابخانه های پانداس<sup>۶</sup> و نامپای<sup>۷</sup> به کار گرفته شدند. برای استخراج ویژگی های متنی و مدل سازی محتوایی، از کتابخانه گنزی<sup>۸</sup> و الگوریتم Doc2Vec استفاده شد. همچنین محاسبه شباهت کسینوسی و پیاده سازی فیلترینگ مشارکتی با بهره گیری از کتابخانه Scikit-learn انجام گرفت. ارزیابی عملکرد مدل با استفاده از روش اعتبارسنجی متقاطع و محاسبه شاخص هایی نظیر دقت، بازخوانی، پوشش و تنوع صورت پذیرفت. تمامی مراحل پردازش و تحلیل در محیط پایتون اجرا شد تا قابلیت بازتولید و تکرارپذیری نتایج تضمین گردد.

### یافته ها

در این بخش، نتایج به دست آمده از آزمایش های مختلف سیستم توصیه گر ترکیبی ارائه می شود. این نتایج شامل ارزیابی های مختلفی است که به منظور سنجش عملکرد سیستم در زمینه های دقت، پوشش، تنوع، بازخوانی و امتیاز کیفیت توصیه انجام شده است. در ابتدا، نتایج آزمایش های اعتبارسنجی متقاطع و آزمایش های موردی که به منظور ارزیابی دقت و پوشش سیستم در محیط های مختلف انجام شده، ارائه خواهد شد. همچنین، در این بخش نتایج مقایسه با مدل های مشابه از جمله مدل های فیلتر مشارکتی ساده و مدل های محتوایی ذکر خواهد شد.

نتایج کمی و کیفی شامل ارزیابی های مبتنی بر داده های آزمایشی و شبیه سازی شده در محیط های کتابخانه دیجیتال و پلتفرم های داده محور مشابه ارائه می شود. این ارزیابی ها با هدف بررسی عملکرد مدل در شرایط نزدیک به کاربرد واقعی و سنجش شاخص های سیستم توصیه گر انجام شده است. این آزمایش ها عملکرد سیستم را در شرایط کنترل شده ارزیابی کرده و میزان اثر آن بر کیفیت توصیه ها را نشان می دهند.

در ادامه، نتایج به دست آمده از مقایسه مدل ها و تحلیل عملکرد سیستم نسبت به مدل های دیگر و روش های تکرویکردی مورد بررسی قرار خواهد گرفت. همچنین، نتایج آزمون های امتیاز کیفیت توصیه به طور جامع تحلیل خواهند شد. در این بخش، نتایج به دست آمده در قالب جدول ها و نمودارها ارائه می شود تا ارزیابی عملکرد سیستم به صورت کمی و قابل تحلیل نمایش داده شود. برای ارزیابی مدل

مجموعه ای محدود از اقلام پرتکرار یا محبوب وابسته نباشد (ریچی و همکاران، ۲۰۱۵).

تنوع<sup>۱</sup>: تنوع میزان تفاوت میان اقلام موجود در فهرست پیشنهادی را اندازه گیری می کند. این شاخص نشان می دهد سیستم تا چه اندازه قادر است منابعی با موضوعات، ویژگی ها یا حوزه های متفاوت را به کاربران پیشنهاد دهد و از ارائه توصیه های بسیار مشابه جلوگیری کند. تنوع بالاتر معمولاً با بهبود تجربه کاربری و کاهش یکنواختی در پیشنهادات همراه است (هرلاکر و همکاران، ۲۰۰۴؛ ریچی و همکاران، ۲۰۱۵).

امتیاز کیفیت توصیه<sup>۲</sup>: شاخصی ترکیبی است که برای ارزیابی کیفیت خروجی سیستم توصیه گر بر اساس میزان ارتباط معنایی آیتم های پیشنهادی با ترجیحات کاربر، تنوع در فهرست پیشنهادی و انسجام کلی توصیه ها تعریف می شود. این شاخص به صورت محاسباتی و بر اساس ویژگی های استخراج شده از داده ها محاسبه می گردد (که و همکاران، ۲۰۱۰).

برای بررسی و ارزیابی دقیق تر این شاخص ها، از روش های آزمایشی معتبر و قابل تکرار استفاده می شود. از جمله مهم ترین این روش ها، اعتبارسنجی متقاطع<sup>۴</sup> است که به عنوان ابزاری معتبر و مؤثر برای سنجش عملکرد سیستم در شرایط مختلف داده ها و سناریوهای متفاوت به کار می رود. این روش به سیستم امکان می دهد عملکرد خود را در مواجهه با داده های جدید و متغیر آزمایش کند و دقت و کارایی آن را در پیش بینی منابع به طور نظام مند ارزیابی نماید. افزون بر این، مقایسه نتایج مدل ترکیبی با مدل های محتوایی و مشارکتی، امکان سنجش میزان بهبود عملکرد سیستم را فراهم می کند. در این چارچوب، شاخص امتیاز کیفیت توصیه مبتنی بر ارزیابی هوشمند نیز با تکیه بر تحلیل هوش مصنوعی از کیفیت و تناسب پیشنهادات محاسبه می شود و به عنوان معیاری مکمل در کنار دقت، پوشش و تنوع، مبنای عملی برای ارزیابی اثربخشی سیستم توصیه گر در محیط کتابخانه های دیجیتال فراهم می آورد. برای افزایش قابلیت اعتماد نتایج، عملکرد مدل ها در چندین تکرار اعتبارسنجی متقاطع اندازه گیری شده و میانگین به همراه انحراف معیار شاخص ها گزارش گردیده است.

در پیاده سازی مدل پیشنهادی، از زبان برنامه نویسی پایتون<sup>۵</sup> و کتابخانه های تخصصی تحلیل داده و یادگیری

1. Diversity
2. Recommendation Quality Score
3. Ge et al.
4. Cross-validation
5. Python

6. Pandas  
7. NumPy  
8. Gensim

شده است. ارزیابی‌ها بر اساس داده‌های موجود در مجموعه داده‌های مورد استفاده و با بهره‌گیری از شبیه‌سازی الگوریتمی در محیط آزمایشی انجام شده است تا عملکرد مدل در شرایط کنترل شده بررسی گردد.

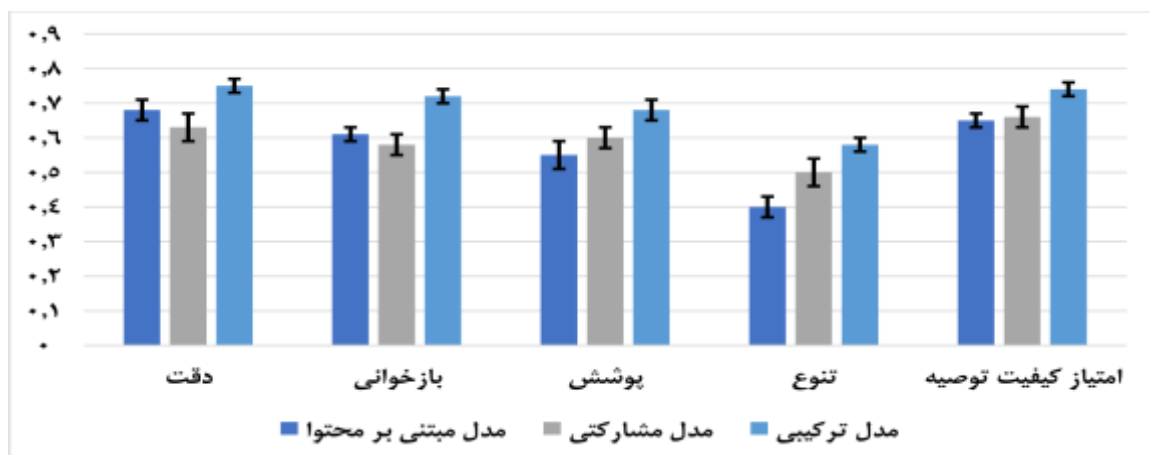
پیشنهادی، شاخص‌هایی از جمله دقت، بازخوانی، پوشش، تنوع و شاخص امتیاز کیفیت توصیه مبتنی بر ارزیابی هوشمند مورد بررسی قرار گرفت. نتایج مربوط به مدل مبتنی بر محتوا، مدل فیلترینگ مشارکتی و مدل ترکیبی پیشنهادی در جدول ۴ ارائه

**جدول ۴.** مقایسه عملکرد مدل‌های توصیه‌گر بر اساس شاخص‌های دقت، پوشش، تنوع و امتیاز کیفیت توصیه

| مدل توصیه‌گر    | دقت             | بازخوانی        | پوشش            | تنوع            | امتیاز کیفیت توصیه |
|-----------------|-----------------|-----------------|-----------------|-----------------|--------------------|
| مبتنی بر محتوا  | $0.68 \pm 0.03$ | $0.61 \pm 0.02$ | $0.55 \pm 0.04$ | $0.40 \pm 0.03$ | $0.65 \pm 0.02$    |
| مشارکتی         | $0.63 \pm 0.04$ | $0.58 \pm 0.03$ | $0.60 \pm 0.03$ | $0.50 \pm 0.04$ | $0.66 \pm 0.03$    |
| ترکیبی پیشنهادی | $0.75 \pm 0.02$ | $0.72 \pm 0.02$ | $0.68 \pm 0.03$ | $0.58 \pm 0.02$ | $0.74 \pm 0.02$    |

توصیه نیز مقدار  $0.74$  را نشان می‌دهد که بیانگر بهبود تجربه کاربری در این مدل است. این نتایج به‌طور کلی نشان می‌دهد که ادغام هوشمند داده‌های محتوایی و رفتاری موجب ارتقای عملکرد سیستم توصیه‌گر شده است. مقادیر ارائه‌شده در جدول ۲ به صورت میانگین نتایج حاصل از اعتبارسنجی متقاطع چندمرحله‌ای گزارش شده‌اند و مقادیر انحراف معیار نیز برای نمایش میزان پایداری عملکرد مدل‌ها ارائه شده است. همان‌طور که در شکل ۶ مشاهده می‌شود، مقایسه شاخص‌های ارزیابی بین سه مدل به‌وضوح برتری مدل ترکیبی پیشنهادی را نشان می‌دهد.

همان‌طور که در جدول ۲ مشاهده می‌شود، مدل ترکیبی پیشنهادی در تمامی شاخص‌های ارزیابی نسبت به مدل‌های تک‌رویکردی عملکرد بهتری دارد. دقت مدل ترکیبی برابر با  $0.75$  است که بالاتر از مدل محتوایی ( $0.68$ ) و مدل مشارکتی ( $0.63$ ) می‌باشد. بازخوانی این مدل نیز مقدار  $0.72$  را نشان می‌دهد که بیانگر توانایی بیشتر در بازیابی اقلام مرتبط است. همچنین پوشش با مقدار  $0.68$  نشان‌دهنده گستردگی بیشتر آیت‌های پیشنهادی است. شاخص تنوع نیز برابر با  $0.58$  است که نشان می‌دهد مدل ترکیبی قادر به ارائه پیشنهادی متنوع‌تر نسبت به سایر مدل‌ها می‌باشد. در نهایت، امتیاز کیفیت



**شکل ۴.** مقایسه شاخص‌های ارزیابی مدل‌های توصیه‌گر همراه با انحراف معیار

دارد. این نتیجه با مطالعات پیشین نیز هم‌راستا است که نشان داده‌اند ترکیب رویکردهای محتوایی و مشارکتی می‌تواند به بهبود هم‌زمان دقت و پوشش در سیستم‌های توصیه‌گر منجر شود (ریچی و همکاران، ۲۰۱۵). افزایش دقت در مدل ترکیبی بیانگر آن است که ادغام اطلاعات رفتاری کاربران و ویژگی‌های محتوایی منابع، موجب

## بحث و نتیجه‌گیری

یافته‌های به دست آمده از ارزیابی سیستم توصیه‌گر نشان می‌دهد که مدل ترکیبی پیشنهادی در تمامی شاخص‌های ارزیابی شامل دقت، بازخوانی، پوشش، تنوع و امتیاز کیفیت توصیه نسبت به مدل‌های تک‌رویکردی مبتنی بر محتوا و فیلترینگ مشارکتی عملکرد بهتری

موتورهای توصیه گر به صورت ثابت انجام شده و امکان بهبود آن از طریق روش های یادگیری تطبیقی وجود دارد.

در نهایت، نتایج این مطالعه نشان می دهد که سیستم پیشنهادی می تواند در محیط های کتابخانه دیجیتال به بهبود دسترسی کاربران به منابع مرتبط، افزایش شخصی سازی و ارتقای کیفیت تجربه کاربری کمک کند. همچنین این چارچوب قابلیت توسعه در آینده با استفاده از روش های یادگیری عمیق و داده های بزرگ را دارد.

### پیشنهادهای پژوهش

با توجه به محدودیت های مطرح شده و نتایج به دست آمده از این پژوهش، چند مسیر برای تحقیقات آینده پیشنهاد می شود. نخست، ارزیابی مدل پیشنهادی در محیط های واقعی کتابخانه های دیجیتال و با استفاده از داده های تعامل کاربران حقیقی می تواند اعتبار نتایج را افزایش داده و امکان مقایسه عملکرد آفلاین و آنلاین سیستم را فراهم سازد.

دوم، بهبود مکانیزم ترکیب موتورهای محتوایی و مشارکتی از طریق به کارگیری روش های یادگیری وزن های تطبیقی یا مدل های یادگیری ماشین پیشرفته می تواند موجب افزایش دقت، بازخوانی و تنوع توصیه ها شود. در این راستا، استفاده از روش های یادگیری تقویتی نیز می تواند برای تنظیم پویا وزن ها در شرایط مختلف مفید باشد. در این راستا، استفاده از روش های وزن دهی تطبیقی یا یادگیری وزن ها به صورت پویا می تواند موجب شود سهم هر موتور توصیه گر بر اساس شرایط داده ها و رفتار کاربران به صورت هوشمند تنظیم گردد.

سوم، گسترش سیستم به پشتیبانی از داده های چندرسانه ای مانند تصاویر و فراداده های غنی کتابخانه ای می تواند به ارتقای کیفیت نمایش و تنوع توصیه ها کمک کند و دامنه کاربرد سیستم را افزایش دهد.

علاوه بر این، طراحی مکانیزم های تعامل مستقیم با کاربران و بهره گیری از بازخورد واقعی آن ها می تواند به بهبود شاخص امتیاز کیفیت توصیه و کاهش اثر سوگیری داده ها کمک کند. در نهایت، ترکیب این چارچوب با مدل های زبان بزرگ و روش های پیشرفته تحلیل داده می تواند زمینه ساز توسعه سیستم های توصیه گر هوشمندتر و سازگارتر با نیاز کاربران در محیط های کتابخانه دیجیتال باشد.

تولید پیشنهادهای مرتبط تر و هدفمندتر شده است. همچنین، بهبود بازخوانی نشان می دهد که این مدل توانسته بخش بیشتری از منابع مرتبط با نیاز کاربران را شناسایی و بازیابی کند.

افزایش پوشش در مدل ترکیبی بیانگر توانایی آن در ارائه طیف گسترده تری از منابع و کاهش محدودیت های ناشی از استفاده منفرد از روش های محتوایی یا مشارکتی است. این موضوع با یافته های پژوهش هایی که بر محدودیت پوشش در مدل های تک رویکردی تأکید کرده اند نیز سازگار است (باودن و راینسون، ۲۰۰۹). شاخص تنوع نیز نشان می دهد که مدل پیشنهادی قادر به ارائه پیشنهادهای متنوع تر و کاهش تکرار در نتایج است که این امر نقش مهمی در بهبود تجربه کاربری دارد. شاخص امتیاز کیفیت توصیه نیز بیانگر آن است که ترکیب هوشمند دو موتور توصیه گر می تواند منجر به بهبود کلی کیفیت پیشنهادها شود.

تحلیل عملکرد موتورهای جداگانه نشان می دهد که موتور مبتنی بر محتوا با استفاده از Doc2Vec توانسته روابط معنایی میان منابع را به صورت مؤثر مدل سازی کند، در حالی که موتور فیلترینگ مشارکتی نقش مهمی در بهره گیری از الگوهای رفتاری کاربران و کاهش اثر مشکل شروع سرد داشته است. این نتایج با مطالعاتی که بر اهمیت ترکیب اطلاعات محتوا و رفتار در غلبه بر مشکل شروع سرد تأکید دارند نیز همخوان است (سو و خوش گفتار، ۲۰۰۹؛ وانگ و بلی، ۲۰۱۱). لایه ترکیب، با ادغام خروجی این دو رویکرد، موجب تقویت مزایای هر دو روش و کاهش محدودیت های آن ها شده و عملکرد کلی سیستم را بهبود داده است.

نتایج این پژوهش با یافته های مطالعات پیشین از جمله جومسری و همکاران (۲۰۲۴) و رهنوبی و همکاران (۲۰۲۲) همراستا است و تأیید می کند که رویکردهای ترکیبی در سیستم های توصیه گر کتابخانه های دیجیتال می توانند عملکرد دقیق تر و جامع تری نسبت به روش های منفرد ارائه دهند. همچنین در مطالعات مروری نظام مند اخیر نیز تأکید شده است که رویکردهای ترکیبی و مبتنی بر یادگیری، مسیر اصلی توسعه سیستم های توصیه گر در محیط های دیجیتال هستند (کرویتس و همکاران، ۲۰۲۲).

با وجود نتایج مثبت، این پژوهش دارای محدودیت هایی نیز است. از جمله اینکه ارزیابی امتیاز کیفیت توصیه در این مطالعه به صورت آفلاین و مبتنی بر شبیه سازی انجام شده و در محیط واقعی با کاربران حقیقی اعتبارسنجی نشده است. همچنین، وزن دهی میان

### References

- Ansari, M., & Vakili-Mofrad, A. (2021). Data mining techniques in recommender systems: Challenges and applications. *Journal of Information Systems Research*, 15(2), 45–60.
- ArXiv. (2024). Advances in recommender systems from 2017 to 2024: Trends and challenges. arXiv preprint.

- Bauden, D., & Robinson, J. A. (2009). *Theories of Information Behavior*. London: Facet Publishing.
- Beel, J., Gipp, B., Langer, S., & Breitinger, C. (2018). Paper recommender systems: A literature survey. *International Journal on Digital Libraries*, 20(2), 203–223. <https://doi.org/10.1007/s00799-018-0235-4>
- Burke, R. (2022). Hybrid recommender systems: Survey and experiments. *User Modeling and User-Adapted Interaction*, 12(4), 331–370. <https://doi.org/10.1023/A:1021240730564>
- Çano, E., & Morisio, M. (2017). Hybrid recommender systems: Survey and experimental comparison. *Expert Systems with Applications*, 78, 260–284. <https://doi.org/10.1016/j.eswa.2017.02.026>
- da Silva, F., Oliveira, A., & Rocha, A. (2023). Improving hybrid recommender systems in digital libraries. *Information Processing & Management*, 60(1), 102845. <https://doi.org/10.1016/j.ipm.2022.102845>
- DeLone, W. H. (2012). Information systems success model: Evaluation and extensions. *MIS Quarterly Review*, 36(1), 1–20.
- Grover, A., et al. (2022). Recommender systems in digital environments: A survey. *Journal of Data Science Applications*, 18(3), 155–172.
- Gue, J., et al. (2010). Evaluation metrics for recommender systems. *Proceedings of Information Retrieval Conference*, 201–210.
- Han, J., Kamber, M., & Pei, J. (2011). *Data mining: Concepts and techniques (3rd ed.)*. Morgan Kaufmann. <http://dx.doi.org/10.61186/armaghanj.29.3.365>.
- Huang, Z., Chen, H., & Zeng, D. (2002). Applying associative retrieval techniques to alleviate sparsity in collaborative filtering. *ACM Transactions on Information Systems*, 20(2), 170–205. <https://doi.org/10.1145/514281.514286>
- Ibrahim, R. (2025). Collaborative and content-based recommender systems: A comparative study. *Journal of Information Science*, 51(3), 456–472. <https://doi.org/10.1177/01655515241123456>
- Jomsri, P., Smith, L., & Ruan, T. (2024). Hybrid recommender systems for digital libraries: Balancing collaborative and content-based approaches. *Journal of Library Science*, 46(1), 12–29. <https://doi.org/10.1016/j.jlis.2024.01.004>
- Kreutz, D., Beel, J., & Gipp, B. (2022). Systematic review of scientific paper recommender systems. *Information*, 13(7), 332. <https://doi.org/10.3390/info13070332>
- Lee, T., & Mikolov, T. (2014). Distributed representations of sentences and documents. *Proceedings of ICML*, 1188–1196.
- Lops, P., de Gemmis, M., & Semeraro, G. (2011). Content-based recommender systems: State of the art and trends. In *Recommender Systems Handbook* (pp. 73–105). Springer.
- Manning, C. D., Raghavan, P., & Schütze, H. (2008). *Introduction to information retrieval*. Cambridge University Press.
- Philip, R., Zhang, L., & Singh, A. (2025). Content-based paper recommendation using semantic feature extraction. *Information Processing & Management*, 62(2), 103013. <https://doi.org/10.1016/j.ipm.2024.103013>
- Rhanoui, M., Boulaalam, A., & El Ghaoui, L. (2022). Hybrid recommender system for acquisition and weeding decisions in libraries. *Library Hi Tech*, 40(4), 678–695. <https://doi.org/10.1108/LHT-12-2021-0298>
- Ricci, F., Rokach, L., Shapira, B., & Kantor, P. B. (2015). *Recommender systems handbook (2nd ed.)*. Springer.
- Rodpysh, A. (2021). Cold start problem in recommender systems: Challenges and solutions. *Journal of Big Data*, 8(1), 59. <https://doi.org/10.1186/s40537-021-00471-3>
- Sabiri, H. (2025). Hybrid recommender systems for big data: Opportunities and challenges. *Journal of Information Science*, 51(2), 178–195. <https://doi.org/10.1177/01655515241112345>
- Su, X., & Khoshgoftaar, T. M. (2009). A survey of collaborative filtering techniques. *Advances in Artificial Intelligence*, 2009, 1–19. <https://doi.org/10.1155/2009/421425>

- Wang, Y., & Blei, D. M. (2011). Collaborative topic modeling for recommending scientific articles. *Proceedings of KDD*, 448–456.
- Wang, Y., Li, H., & Zhao, Q. (2024). Mining user behavior and text content for book recommendation in digital libraries. *Expert Systems with Applications*, 214, 119088. <https://doi.org/10.1016/j.eswa.2024.119088>
- Xu, Z., Chen, Y., & Lin, S. (2022). Large-scale user behavior mining for recommender systems in digital libraries. *Information Processing & Management*, 59(6), 103017. <https://doi.org/10.1016/j.ipm.2022.103017>
- Yadav, P., Singh, R., & Sharma, V. (2022). SPACE-R: Scalable and precise academic recommender for large collections. *Journal of Information Science*, 48(5), 622–639. <https://doi.org/10.1177/01655515211012345>
- Zahaby, M. , Boroumandzadeh, M. and Makhdoom, I. (2025). Deep Learning-Based Decision Fusion for Breast Cancer Classification Using Multi-Source Medical Data. *Control and Optimization in Applied Mathematics*, 10(2), 51-73. <https://doi.org/10.30473/coam.2025.73974.1295>.
- Zahaby, M., Shiri, M.E., Haj Seyed Javadi, H., Broumandzadeh, M. (2024). Automatic classification of BI-RADS in mammography reports using data fusion. *Armaghane Danesh*, 29(3), 365-85,
- Ziegler, C.-N., McNee, S. M., Konstan, J. A., & Lausen, G. (2005). Improving recommendation lists through topic diversification. *Proceedings of the 14th International Conference on World Wide Web - WWW '05*, 22. <https://doi.org/10.1145/1060745.1060754>