

Research of Cross-Domain Recommendation Algorithm Based on Knowledge Graph

Yancong Zhou, *IAENG, Member*, Shan Jiao, Weiwei Li, Kaiyue Zeng, Dongdong Wang

Abstract—To address the challenge of insufficient interactive data in cross-domain recommendations, this study proposes an innovative approach by integrating the comment text information of similar users, obtained through a Knowledge Graph (KG), into the Recommendation Model (RM) as auxiliary information. Specifically, features are extracted from the comment text, and KG is incorporated into the cross-domain Recommendation System (RS) to capture richer semantic features. A novel method for identifying similar users is introduced, in which KG paths are semantically represented using the TransD model. To calculate user similarity and generate a more accurate user similarity matrix, a user-based collaborative filtering algorithm is employed. This approach enhances the acquisition of user information by improving the matrix's accuracy. Furthermore, the inclusion of auxiliary information increases the diversity of training data, which improves recommendation performance for cold-start users and enhances both the interpretability and accuracy of the recommendations. The proposed end-to-end cross-domain recommendation system spans books, movies, and music by integrating aspect text features and KG-based auxiliary information. The model was evaluated on the Amazon dataset using MAE and MSE as performance indicators. On the book-movie dataset, the model's performance improved by 2.66% (MAE) and 4.8% (MSE) compared to the best-performing baseline. Similarly, on the book-music dataset, the model outperformed the best baseline by 3.85% (MAE) and 1.42% (MSE). Experimental results demonstrate that the new model outperforms other comparison models, thereby validating its effectiveness.

Index Terms—Aspect level text analysis, Cross-domain recommendation, Knowledge graph, Comment text

I. INTRODUCTION

With the rapid development of mobile internet, online shopping not only provides more choices for people, but also increases the difficulty of finding the goods. To address the problem of information overload, Recommendation Systems (RS) have emerged. However,

traditional recommendation approaches face challenges when interaction data between users and products is sparse. So In response, some researchers have proposed cross-domain recommendation systems based on transfer learning, which used relatively rich information from source domains to improve the recommendation performance of sparse domains. A key assumption of cross-domain recommendation is that there is consistency or correlation between cross-domain user preferences or item features. This consistency and correlation primarily involve factors such as the overlap between users and items, the similarity of user preferences, similar features between items, and the relationship between latent factors, all of which help compensate for the lack of data in the target domain., so as to make up for the lack of information in the target domain.

Comment text provides valuable insights into the characteristics of both users and products. It can not only alleviate the sparsity of the scoring data, but also enhance the interpretability of the recommendation algorithm. Crucially, it helps establish connections between source and target domains. Most existing cross-domain recommendation models are only based on the scoring data or simply relies on the knowledge sharing and transferring of the comment text, resulting in insufficient exploration of comment information. Fully exploring and making use of the implicit information in the comment text will help alleviate the problems of cold start and sparse data [1]. Moreover, broad textual analysis tends to overlook finer details and fails to capture subtle emotional expressions. Therefore, the various fine-grained views implied in the comment information are significant[2]. In order to conduct a more complete user preference analysis, the system need to discover information on all aspects of the comment and determine the views expressed by each aspect's text.

In addition, most of the cross-domain recommendation models are based on overlapping user data. When there are fewer ones, cross-domain recommendation will be difficult, and the interpretation is not strong[3]. Knowledge Graphs (KG) can efficiently identify entities and relationships between users and items based on users' historical behavior, so RS can obtain more rich user's and item's background information, achieve more accurate and effective recommendation[4], improve the interpretability, enrich the diversity of recommendation, and enhance users' satisfaction. Therefore, KG is integrated into RS as auxiliary information.

The contributions are summarized as follows:

(1) Aspect-level text analysis techniques were employed to extract comment information, allowing for a comprehensive exploration of comment content, and improved the recommendation performance. The Word2Vec method was

Manuscript received May 6, 2024; revised December 26, 2024.

This work was supported in part by Tianjin Graduate Research Innovation Program under Grant 2022SKYZ312.

Yancong Zhou is a professor of Tianjin University of Commerce, Tianjin 300134, China (corresponding author to provide phone:022-26667577 fax:022-26667577, e-mail: zycong78@126.com).

Shan Jiao is a postgraduate student of Tianjin University of Commerce, Tianjin 300134, China (e-mail: j13231943757@163.com).

Weiwei Li is a postgraduate student of Zhejiang University of Technology, Zhejiang, 310023, China (e-mail: lww20201920@163.com).

Kaiyue Zeng is a postgraduate student of Tianjin University of Commerce, Tianjin 300134, China (e-mail: 46448526@qq.com).

Dongdong Wang is a postgraduate student of Tianjin University of Commerce, Tianjin 300134, China (e-mail: wangdongdong0127@163.com).

used to obtain word vector representations, and the convolutional layer to capture the context information around each word. Then, used the gating mechanism to identify contextual features related to aspects, and an attention mechanism highlighted the key aspects to derive the final aspect feature representation.

(2) A new method for acquiring similar users was proposed. The TransD model was used to represent the map path semantically, and the user-based collaborative filtering algorithm was integrated to calculate the user similarity, so as to capture the user information more fully and obtain a more accurate similar user matrix.

(3) Cross-domain recommendation was achieved by integrating aspect-level text features and KG auxiliary information. Firstly, fine-grained text analysis data was used to fully mine the user's and item's features. Then, KG as an auxiliary information was integrated into the model to fully explore similar users with its rich semantic path.

(4) The experiment was carried out on the Amazon book, movie and music data set, and a good performance was obtained. On the book-movie dataset, performance improved by 2.66% (MAE) and 4.8% (MSE) compared to the best-performing model. On the book-music dataset, MAE was reduced by 3.85% and 1.42%, respectively.

II. RELATED WORK

A. Cross-domain Recommendation

In order to alleviate the user cold start problem, cross-domain recommendation systems (RS) were developed. It used richer information in the source domain to improve the recommendation performance of the target domain. Early research primarily focused on the traditional collaborative filtering methods, which had high requirements for user-item interaction data. With the development of deep learning, researchers began to learn cross-domain recommendations based on deep learning. Zhu et al.[5] used the sparsity of individual user and item in source domain and target domain to guide the training process of fully connected neural networks to make full use of the scoring data. Wang et al.[6] learned more accurate potential features of users in sparse domains and used a neighbor-based feature mapping method to learn cross-domain potential feature mapping functions. The above models were mostly based on user ratings, labels and other data, they largely ignored the comment data.

For cold-start users, representation mapping is performed based on models trained using overlapping users. However, other issues remain, such as insufficient interactive data, difficulty in explicitly distinguishing fine-grained semantic features, and the accumulation and amplification of noise information during intermediate steps[7]. By introducing auxiliary information, interactive information can be enriched. Based on the idea of hybrid recommendation algorithm, Lu[8] integrated scores, users and items related information and implicit feedback information in a joint training way to conduct cross-domain recommendation. Xiao[9] proposed a model based on dual noise reduction autoencoder, that fused auxiliary information and scoring data to extract features of users and items.

B. Comment Text Recommendation

Recent research has increasingly utilized comment text to learn the representation of users and items. For example, DeepCoNN [10] used two parallel networks to integrate the comment information into documents for learning. Seo et al.[11] proposed the D-Attn model, which used a document-level approach to model the initial representation of users and items. But static encoding forms struggled to capture the user's preference. In order to learn the dynamic representation of target users for target items, Chin et al.[12] proposed an ANR model. It composed the items' and users' comment into a document and mapped the words into continuous vector representation. Then transformed the representation into different aspects through the defined K aspect matrix. Finally score prediction was given using aspect representation and aspect importance..

Regarding comment-based approaches, Liu et al.[13] believed the same words or similar comments may contain different information for different users and items, so they integrated the ID information into the construction of NRPA model, and FM was used for score prediction. Tay et al.[14] proposed the MPCN model. The main contribution was to use Co-Attention at the comment level and word level.

To sum up, most of the existing comment-based methods are coarse-grained information extraction of comments, and the information of the comments has not been fully mined. So the mining and utilization of fine-grained text information is the focus of current research on text recommendation.

C. Recommendation Based on KG

The Knowledge Graph (KG) is a heterogeneous graph structure that contains rich semantic information, which can significantly enhance model performance. Path-based approaches in KGs use the connections between items to make recommendations, which need to measure the similarity of the connection between entities. Cui[15] proposed HeroGRAP, which constructed shared graphs by collecting users and items, modeled the graph structure to obtain information about each domain, and introduced a circular attention mechanism to seek neighbor nodes. In the embedded-based approaches, the relations and nodes in KG are represented by vectors in the low-dimensional space, and the recommendation is assisted according to the obtained embedding vectors. Ye et al.[16] proposed the BEM model, which used Bayesian model to adjust the vector representation of users and items obtained through TransE model and graph neural network, in which the two kinds of graph structure data of KG and behavior graph were used. Since the above methods need to learn the structured knowledge vector representation first, end-to-end recommendation is challenging. As a result, multi-task learning models have garnered increased attention. Cao et al.[17] proposed KTUP to represent the knowledge of the preference relation between users and items through TransH, obtain the feature vector, complete KG, and enrich the vector representation of item preferences.

Through the analysis of the above studies, it can be seen that KG-assisted RS can enrich the recommendation data, complete KG path, and improve the performance. Therefore, to achieve good performance the KG can be introduced into RS, especially when the data is sparse.

III. METHODOLOGY

A. Aspect Level Text Feature Extraction

Aspect word extraction refers to extracting all comment entities or attributes of entities in the comment. Previous studies based on text analysis are mostly based on the level of word or sentence, which are not conducive to the full capture of comment information. An aspect represents a higher-level semantic attribute, and it is an attribute of the item expressed through the comment by the users. Aspect-level text analysis technology can capture text information at a finer granularity.

(1) Comment Text Embedding

Word embedding technology can embed a high-dimensional vector space into a low-dimensional vector space, and each word or phrase is mapped to a vector over the field of real numbers. This approach reduces the number of parameters, which speeds up model training, and easy to compare, generalize and transfer knowledge. Here, Word2Vec method is used to represent the features of user and comment text. Word vector is used to describe the semantic information of words, and obtain the embedding matrix. The Skip-gram is selected to obtain context information which is used to help better capture the existing comment information.

Given a user document $D_u = [w_1, w_2, \dots, w_l]$, project each word into its embedded representation: $E_u = [e_1, e_2, \dots, e_l]$, $e_j \in R_d$, where l is the document length and d is the word embedded dimension. To capture context information around each word, perform the convolution operation using ReLU(Rectified Linear Unit). By performing CNN on the matrix E_u , the resulting eigenmatrix is $C_u = [c_{1,u}, c_{2,u}, \dots, c_{l,u}]$, where $c_{j,u} \in R_n$ is the potential context eigenvector of the j th word.

(2) Aspect Feature Extraction Based On GLU Gating Mechanism

The GLU (Gated Linear Unit) gating mechanism is employed to identify features relevant to various aspects. A gating mechanism is adept at sensing contextual information, resulting in more accurate feature representations and enhancing the ability to combine features effectively. It can learn distinct feature distributions and, with the same input, control the outflow of information through the gating network, providing different inputs to parallel sub-networks. This allows each sub-network to learn different feature distributions [18].

GTU (Gated Tanh Unit) and GLU are two commonly used gating mechanisms in text analysis[19]. Experimental results by N. Delphin et al. [20] show that GLU achieves faster convergence and higher accuracy compared to GTU. Additionally, GLU helps mitigate the "vanishing gradient" problem in deep architectures by offering linear propagation paths for gradients, while maintaining certain nonlinear capabilities. Therefore, we select the GLU model for aspect feature extraction in this study..

After obtain the context feature vector for aspect m , the aspect feature of the extracted word is shown in formula (1):

$$g_{m,j,u} = (W_m c_{j,u} + b_m) \odot \sigma(W_m^g c_{j,u} + b_m^g) \quad (1)$$

Where σ is the sigmoid activation function and $\odot$ is the element-by-element product operation. $W_m, W_m^g \in R^{k \times n}$, $b_m, b_m^g \in R^k$ represent the transformation matrix and

bias vector of the aspect m respectively. $c_{j,u} \in R_n$ is the potential context feature vector for the j th word. The user document under aspect m can be expressed as formula (2):

$$G_{m,u} = [g_{m,1,u}, g_{m,2,u}, \dots, g_{m,l,u}] \quad (2)$$

Through the above process, M aspect specific word context features G_u can be obtained, as shown in formula (3), and these words can be used for further aspect extraction.

$$G_u = [G_{1,u}, G_{2,u}, \dots, G_{M,u}] \quad (3)$$

(3) Aspect Characteristic Attention Mechanism

Comments from different fields often emphasize different aspects of the items being reviewed[2]. The attention mechanism enables the system to focus on extracting the most relevant and significant information from the input data[21]. In this context, we introduce the attention mechanism to selectively prioritize important features, thereby optimizing the use of available information and improving the accuracy of the recommendation. Two global aspect feature matrices $V_s = [v_{1,s}, v_{2,s}, \dots, v_{M,s}]$ and $V_t = [v_{1,t}, v_{2,t}, \dots, v_{M,t}]$ are used to guide aspect extraction in the field of books and movies. Specifically, aspect m extracted by $G_{m,u}$ is represented as $\delta_{m,u}$, as shown in formula (4) :

$$\delta_{m,u} = \sum_{j=1}^l \theta_{m,j,u} g_{m,j,u} \quad (4)$$

$$\theta_{m,j,u} = \frac{\exp(g_{m,j,u}^T v_{m,s})}{\sum_{i=1}^l (g_{m,i,u}^T v_{m,s})} \quad (5)$$

$\theta_{m,j,u}$ indicates the importance of the word w_j to each aspect, see formula (5). $g_{m,j,u}$ represents the aspect features of the extracted word w_j . So the representations of M aspects can be obtained by D_u and the aspect matrix $A_u = [a_{1,u}, a_{2,u}, \dots, a_{M,u}]$ is formed. Follow the same steps to extract M aspects from D_i : $A_i = [a_{1,i}, a_{2,i}, \dots, a_{M,i}]$. The D_u and D_i aspect extraction parameters are shared in each learning flow. A different set of parameters is used in each learning flow. V_s and V_t are shared in their respective corresponding domains. The Aspect feature acquisition model is as fig.1.

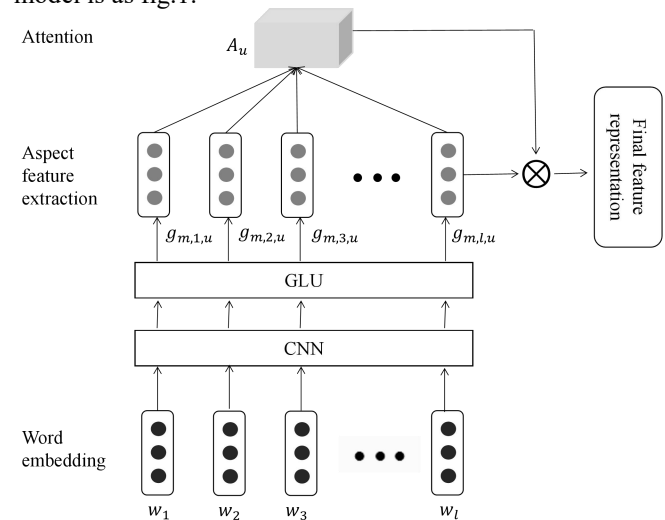


Fig. 1. Aspect feature acquisition model

B. Construction of Book KG

(1) KG Construction

The Knowledge Graph (KG) is a critical component of

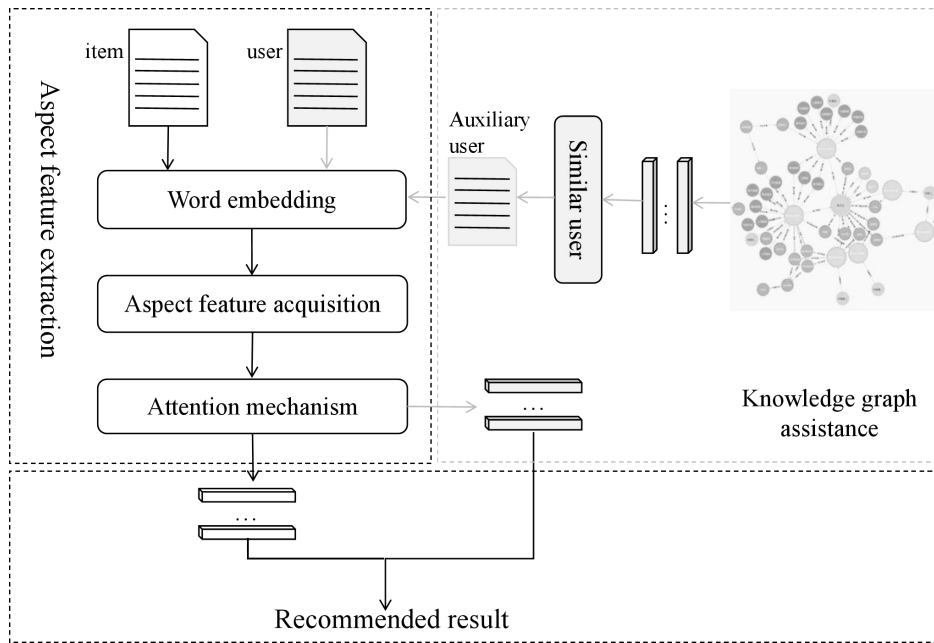


Fig. 2. Cross-domain RM based on KG.

knowledge engineering, used to represent concepts in the physical world and their interrelations in symbolic form. In this study, knowledge extraction was performed on an open dataset from Amazon Books.. The first step involved extracting entities from the dataset, followed by identifying the relationships between these entities. This process enabled the creation of the dataset and the serialization of data objects. Nine entity tags were defined for the construction of the KG: bookid, book_type, helpful, overall (rating), reviewText, reviewTime, reviewerID, reviewerName and summary. To capture user preferences, data with low ratings (overall<4) were excluded, and only items with higher ratings were considered as users' favorite books. A total of 870 data records were selected for KG construction, with 4598 node labels and 4617 relation types identified.

(2) KG Storage and Visualization

Given the specialized nature of the data, the KG was stored in a graph database—Neo4j. This storage method ensures that the increasing complexity of data and its relationships does not degrade search efficiency [22]. Using the data stored in KG, paths of different semantics between node entities can be extracted. Then the user's preference for books can be deduced. Based on KG, each book was taken as a bridge to find the users associated with the target user. When the inference path from user to book was extracted, because of the high complexity of KG, there were often multiple inference paths between the connecting users and the book entity. Sun Y et al.[23] showed that shorter inference paths are more important for recommendation results, and paths longer than five steps tend to introduce noise. To speed up the inference process and model training, only the shortest inference paths were considered. For instance, a maximum of five paths were extracted between any user and a book. The shortest path is chosen for each user-book pair. For example, the inference path "User1 → Book1 → Type A ← Book2" indicates that user1 had read book1, book1 belongs to type A, and book2 also belongs to type A, meaning user1 is likely to be interested in book2. Multiple semantic paths were then integrated to construct the final book KG.

C. Cross-domain Recommendations Fused with Fine-grained Comment Text Based on KG

(1) Model Introduction

The proposed cross-domain recommendation model based on KG is depicted in Fig. 2. In this model, fine-grained user and item features are extracted using aspect-level text analysis technology. The book domain KG is constructed, and knowledge representation learning techniques are used to obtain informative features from the KG. These KG features are then integrated with the aspect-level user features to provide enriched user-item interaction data, which is crucial for cross-domain recommendations, such as recommending books to movies. This fusion of text-based and KG-based features aims to improve recommendation accuracy. The model consists of three key components: aspect level text feature extraction, KG auxiliary information acquisition, feature fusion and cross-domain recommendation.

(2) KG Represents Learning

Many entities in KG have only a little relation connected, which makes it difficult to recognize and reason effectively. Knowledge Representation Learning (KRL) maps semantic information of entities or relation into low-dimensional dense real-valued vector spaces, which makes semantically similar objects close to each other. This approach brings semantically similar objects closer together, which helps alleviate issues like low computational efficiency and sparse data. The KG contains rich semantic information that can be leveraged to obtain more detailed user and item features through KG paths.

For example, in Fig. 3, two paths describe the relationships between books of the same type. These paths suggest that "Gone with the Wind" might be a book a user could be interested in, based on the type similarity with other books the user has interacted with. Similarly, Fig. 4 illustrates additional semantic paths derived from the constructed KG that further enhance the recommendation process.

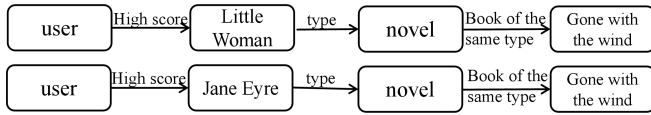


Fig. 3. KG Path example

Through this path relation, more relationships and data information can be captured. By using simple knowledge reasoning, more semantic matching and correlation between multiple entities or features can be identified. The attribute information and relationships of users can be introduced as important auxiliary information, addressing the issue of data sparsity and reducing the cost of acquiring and storing large amounts of auxiliary data. Furthermore, the use of KG enhances recommendation interpretability and enables a deeper discovery of users' interests.. Book recommendation based on KG can discover the heterogeneous relation between books[24].

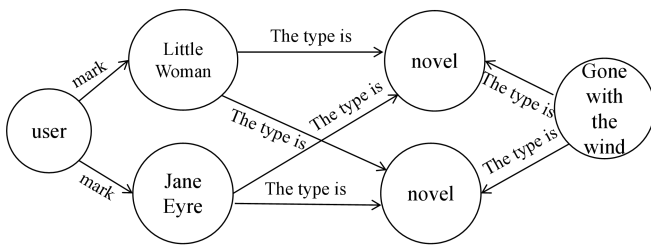


Fig. 4. Path relation

The TransD is introduced, embedding knowledge representation into a continuous vector space. In TransD, each named symbolic object (entity and relation) is represented by two vectors. The first vector captures the meaning of the entity (or relation), while the second vector is used to construct the mapping matrix. For instance, given a triplet $\langle h, r, t \rangle$, its vectors have h, h_p, r, r_p, t, t_p , where the subscript p represents the projection vector, $h, h_p, r, r_p, t, t_p \in \mathbb{R}^n$ and $t, t_p \in \mathbb{R}^m$.

For each triplet $\langle h, r, t \rangle$, set up two mapping matrices $M_{rh}, M_{rt} \in \mathbb{R}^{m \times n}$ to project the entities from the entity space into the relational space. They are as formula (6),(7) :

$$M_{rh} = r_p h_p^T + I_{m \times n} \quad (6)$$

$$M_{rt} = r_p t_p^T + I_{m \times n} \quad (7)$$

Thus, the mapping matrix is determined by entities and relations together. The operation makes the two projection vectors fully interact. When initialize each mapping matrix with the identity matrix, add $I_{m \times n}$ to $M_{r,h}$ and $M_{r,t}$. Define the projection vector, as shown in formula (8):

$$h_{\perp} = M_r^T h \quad t_{\perp} = M_r^T t \quad (8)$$

The scoring function is shown in formula (9) :

$$f_r(h, t) = -\|h_{\perp} + r - t_{\perp}\|_2^2 \quad (9)$$

Use TransD to train all triples, and represent entity i as a vector e_i . The embedded representation of path m in the inference path formed by a pair of entities user-item is shown in formula (10), Where $e_{u,i,m,n}$ represents the embedding vector of entity n in path m from user u to book i .

$$p_{u,i,m} = [e_{u,i,m,1}, e_{u,i,m,2}, \dots, e_{u,i,m,n}] \quad (10)$$

(3) Similar User Acquisition Based on KG

TransD[25] is used for KG representation learning. First,

the low-dimensional vector representation of user features is obtained through TransD, and the semantic similarity between users is calculated to identify the user's semantic nearest neighbors. Secondly, the user behavior similarity matrix sim_{cf} is obtained from the user-book scoring matrix by collaborative filtering recommendation algorithm; Finally, by setting the appropriate fusion ratio, the two are fused to generate the similarity user matrix. The aspect feature acquisition model is shown in Fig. 5:

After constructing and embedding the KG, user similarity is calculated. The user i and j adopt Euclidean distance of the same normal form to calculate the similarity, as shown in formula (11) :

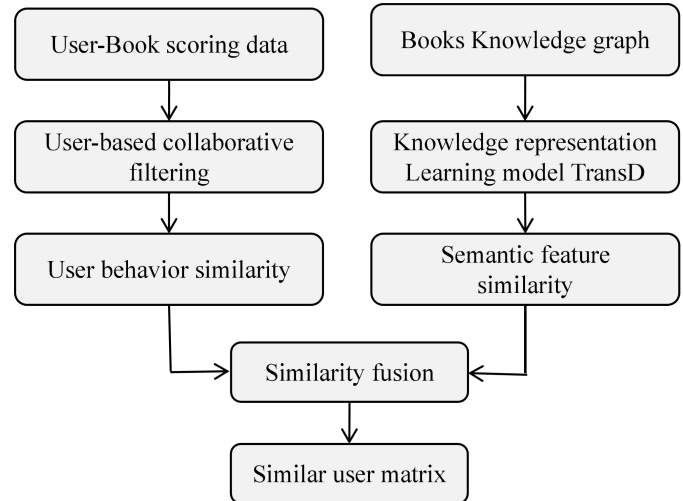


Fig. 5. Similar user acquisition model

$$sim_{trans}(I_i, I_j) = 1 - \frac{\|I_i - I_j\|}{\|I_i - I_j\| + 1} = \frac{1}{\|I_i - I_j\| + 1} \quad (11)$$

The behavioral similarity sim_{cf} and semantic similarity sim_{trans} are integrated using a linear weighting method. The fusion formula is shown in formula (12)

$$sim(i, j) = (1 - \alpha)sim_{cf}(I_i, I_j) + \alpha sim_{trans}(I_i, I_j) \quad (12)$$

User similarity matrix S can be obtained, and similar users can be selected from S for the selection of auxiliary information to enrich cross-domain recommendation data.

(4) Feature Fusion and Cross-domain Recommendation

In cross-domain recommendations, the number of overlapping users across domains is usually very limited. The data sparsity problem is further exacerbated by the scarcity of reviews[26]. The feature vector obtained from KG is fused with the user and item vector to obtain the final user and item characteristics. The comment information of users similar to the interactive users is used as auxiliary information to increase the interactive users information from source domain to target domain, thus alleviates the problem of sparse interactive data in cross-domain recommendation.

The auxiliary documents are formed by different users, which have different language styles and preference priorities for target users. Directly combining the original features of users with the constructed KG feature vectors may lead to incompatibility issues. Wu et al.[26] proposed that placing CNN on top of the context matrix was effective for rating prediction, especially when the semantics were inconsistent. Therefore, after the establishment of user-assisted documents,

another convolutional layer was superimposed between the convolutional layer and the gating mechanism. Finally, an abstract feature matrix $C_{u_a} = [c_{1,u_a}, c_{2,u_a}, \dots, c_{l,u_a}]$ is formed, where $c_{j,u_a} \in \mathbb{R}^n$.

The aspect matrix A_{u_a} is obtained from D_{u_a} through the process of aspect gating mechanism and aspect attention mechanism. In order to effectively use A_{u_a} to update A_u , a gate mechanism based on the interaction of elements of the corresponding aspect is adopted, as formula (13), (14):

$$g_a = \sigma(W_f^1[(A_u - A_{u_a}) \oplus (A_u \odot A_{u_a}) + b_f^1]) \quad (13)$$

$$A_u = \tanh(W_f^2[A_u \oplus (g_a \odot A_{u_a}) + b_f^2]) \quad (14)$$

Where $\oplus$ is a series operation, $W_f^1, W_f^2 \in \mathbb{R}^{k \times 2k}$ are transformation matrix, b_f^1, b_f^2 are bias vectors. The aspect representation A_u is updated to better describe user u .

The final aspect characteristics A_u and A_i of user u and item i are obtained. For a particular user-item pair, the matching score can only reflect the semantic correlation between the two aspects. As not all aspect pairs are equally important, it is beneficial to identify global cross-domain aspect correlations, so important aspect pairs can be highlighted at the domain level to better predict ratings. A set of global aspect representations V_s and V_t can be used to guide aspect extraction for cross-domain preference matching. V_s and V_t can be used to calculate the global cross-domain aspect correlation matrix S , as formula (15):

$$S = \text{LeakyReLU}(V_s^T W V_t) \quad (15)$$

$S(a_1, a_2)$ reflects the importance of shifting preferences based on source domain aspects a_1 and target domain aspects a_2 . $S \in \mathbb{R}^{M \times M}$, $W \in \mathbb{R}^{k \times k}$ are learnable matrix of affinity projection. The LeakyReLU function is adopted to support sparse aspect correlation across domains by setting the corresponding learning rate α to a very small value. It is an activation function specifically designed to solve the Dead ReLU problem, by giving a very small linear component of x to a negative input of $0.01x$ to adjust the negative zero gradient problem. The semantic match between the aspects pairs of A_u and A_i is calculated, as formula (16):

$$S_{u,i} = A_u^T W A_i \quad (16)$$

$S_{u,i}(a_1, a_2)$ reflects the matching degree between corresponding aspects. W is used for affinity projection. Finally, S is used as the attention weight, and pin-aspect matching as the final score prediction. As formula (17), (18):

$$S_{u,i}^r = S \odot S_{u,i} \quad (17)$$

$$\hat{r}_{u,i} = \frac{\sum_{a_1=1}^M \sum_{a_2=1}^M S_{u,i}^r(a_1, a_2) + b_u + b_i}{M^2} \quad (18)$$

Where b_u and b_i represent user bias and item bias.

IV. EXPERIMENTS

A. Datasets

The experiment used the Amazon dataset, a classic dataset in recommend systems (RS). The rating and review information of Books, Movies and TV, and CDs and Vinyl of the 2018 edition, as well as metadata of the Books dataset were used.

First, the data were pre-processed. Interaction records without comments were deleted, and words with high frequency (greater than 0.5) were removed. Random seeds

were used to select half of the overlapping user set as the cold-start user set, with the other half as the non-cold-start user set. The non-cold-start user set was divided into the ordinary user set and the full user set. In addition, a portion of the users in the ordinary user set were randomly selected as the ordinary user set of the training set, and the rest were the non-ordinary user set. The cold-start user set was used as the verification set and the test set. Long documents were truncated (to under 500 words). Words with the highest ranking based on the TD-IDF score were selected as the vocabulary (the top 20,000 words), and the remaining words were removed from the original document. Words were mapped to word vectors using Google's pre-trained word vector model.

B. Experimental Setting

In the experiment, Mean Squared Error (MSE) and Mean Absolute Error (MAE) were used as evaluation metrics. MSE represents the mean squared errors between the predicted and original values. MAE represents the mean absolute errors between the predicted and observed values. Grid search was used to optimize hyperparameters. In the aspect-level text analysis module, the number of aspects was chosen from $\{3, 5, 7, 9\}$, the number of convolution kernels was 150, the potential dimension k was 32, and the window size was 3. The batch size was set to 256. To prevent overfitting, the dropout strategy was applied. The dropout retention probability was set to 0.8, and the initial learning rate was 0.001. The regularization coefficient was set to 1×10^{-5} , the maximum number of training epochs was 200, and the proportion of data used for validation was 0.2. The final performance report was based on more than 5 runs.

C. Performance Comparisons

TABLE I
EXPERIMENTAL RESULTS ON THE BOOK-MUSIC DATASETS SCENARIO.

Models	MAE	MSE
DeepCoNN ^[10]	1.0041	1.2081
RC-DFM ^[28]	0.8997	1.1365
CDLFM ^[6]	0.9102	1.1117
ANR ^[12]	0.9563	1.2231
Our model	0.8651	1.0959

TABLE II
EXPERIMENTAL RESULTS ON THE BOOK-MOVIE DATASETS SCENARIO.

Models	MAE	MSE
DeepCoNN ^[10]	1.0121	1.2216
ALFM ^[27]	0.9945	1.1932
RC-DFM ^[28]	0.8765	1.1412
EMCDR ^[29]	0.8654	1.1293
SSCDR ^[5]	0.8645	1.1212
CDLFM ^[6]	0.9036	1.1264
ANR ^[12]	0.8531	1.1231

The comparison models used were divided into four categories: comment text-based models (represented by DeepCoNN^[10]), aspect-based models (such as ALFM^[27]), cross-domain models (such as RC-DFM^[28], EMCDR^[29], SSCDR^[5], CDLFM^[6]) and aspect-level comment text single-domain models (represented by ANR^[12]). Through comparative experiments, it has been verified the proposed model based on aspect-level text analysis technology and KG assistance can effectively improve the performance.

According to the comparative results in Table I and Table

II, the MAE and MSE were significantly lower than those of other models. On the book-movie dataset, the MAE performance improved by 2.66% compared to ANR[12] (the best-performing model), and the MSE performance improved by 4.8% compared to SSCDR[5] (the best-performing model). On book-music dataset, the MAE and MSE performances improved by 3.85% and 1.42%, respectively, verifying the superiority of the proposed model. Additionally,

the performance of our model and those from literature [3][6][28][29] on the book-movie dataset outperformed models from literature [10] and [27], indicating that cross-domain recommendation effectively alleviates the cold-start problem. The performance comparison The performance comparison on the Book-Music and Book-Movie datasets is shown in Fig. 6 and Fig. 7.

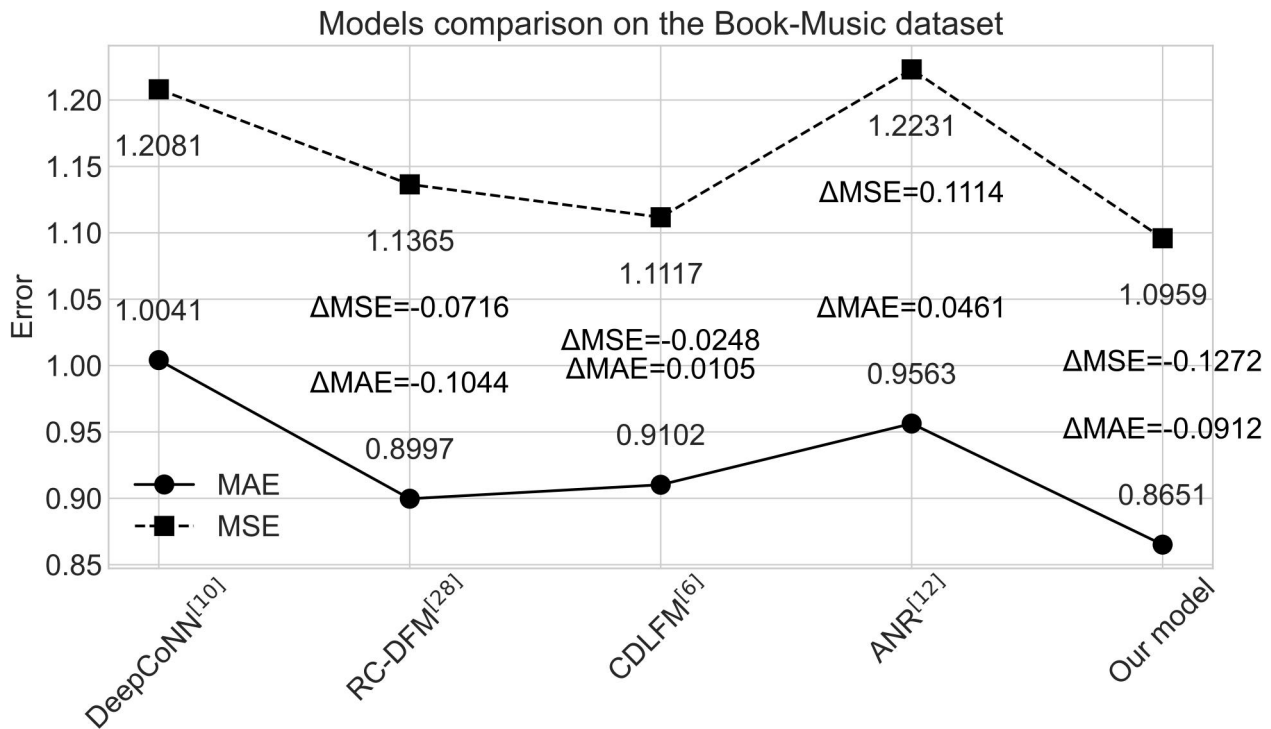


Fig. 6. The comparison on the Book-Music dataset

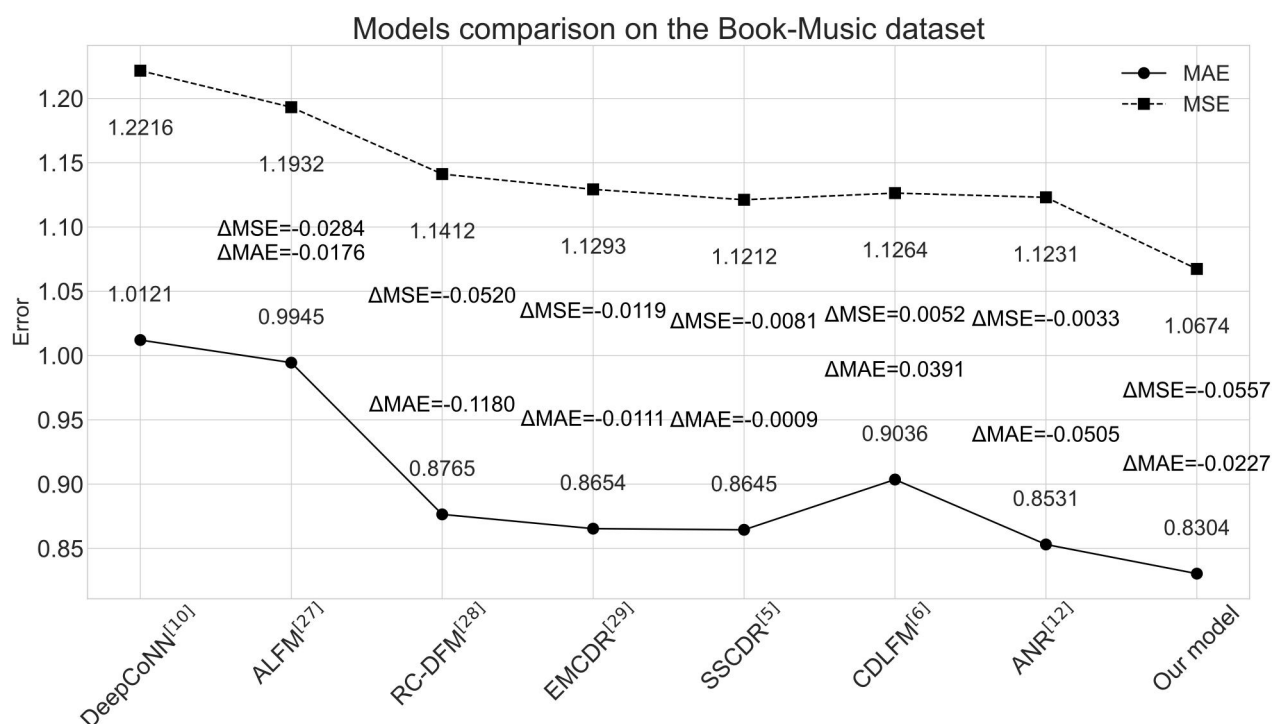


Fig. 7. The comparison on the Book-Movie dataset

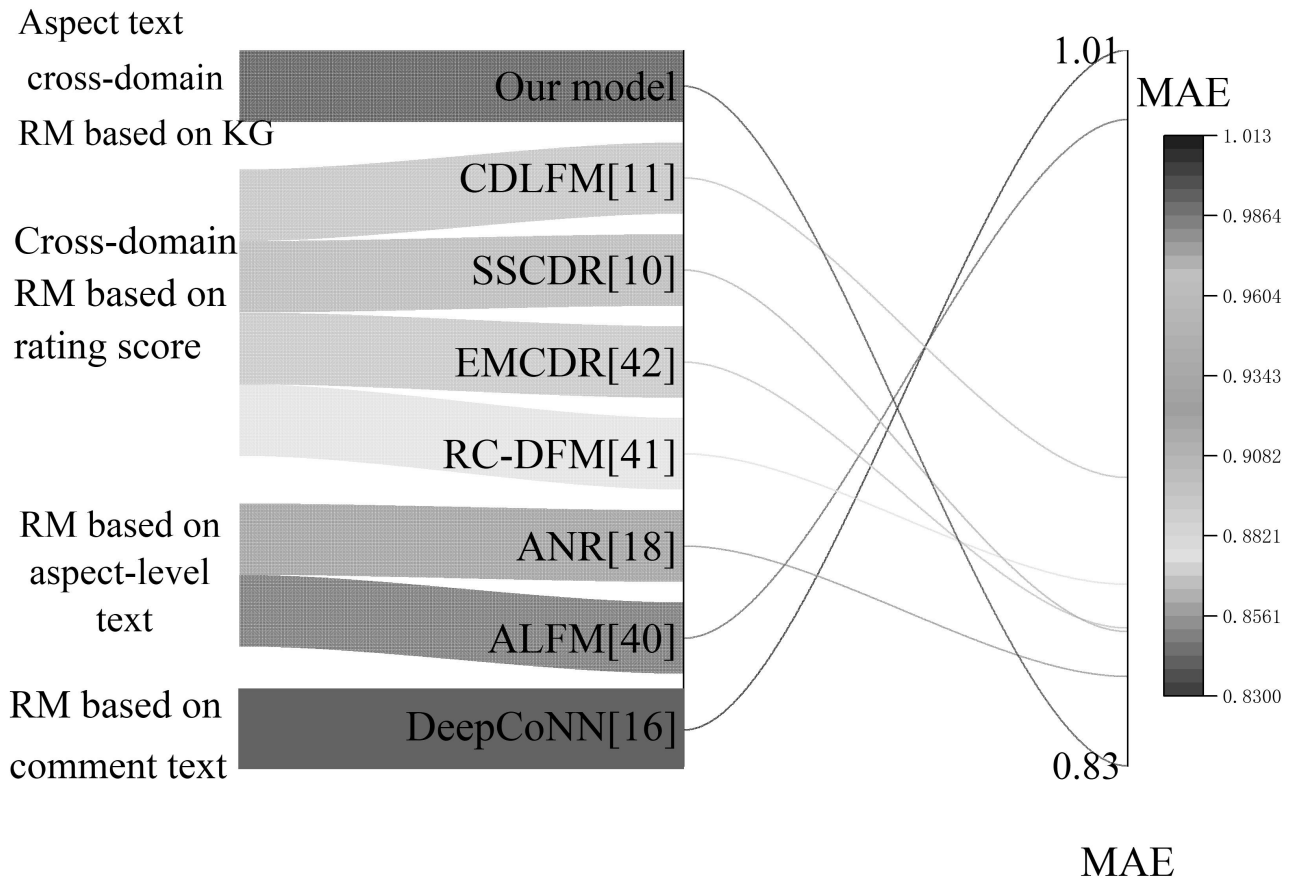


Fig. 8. The MAE comparison on the Book-Music dataset

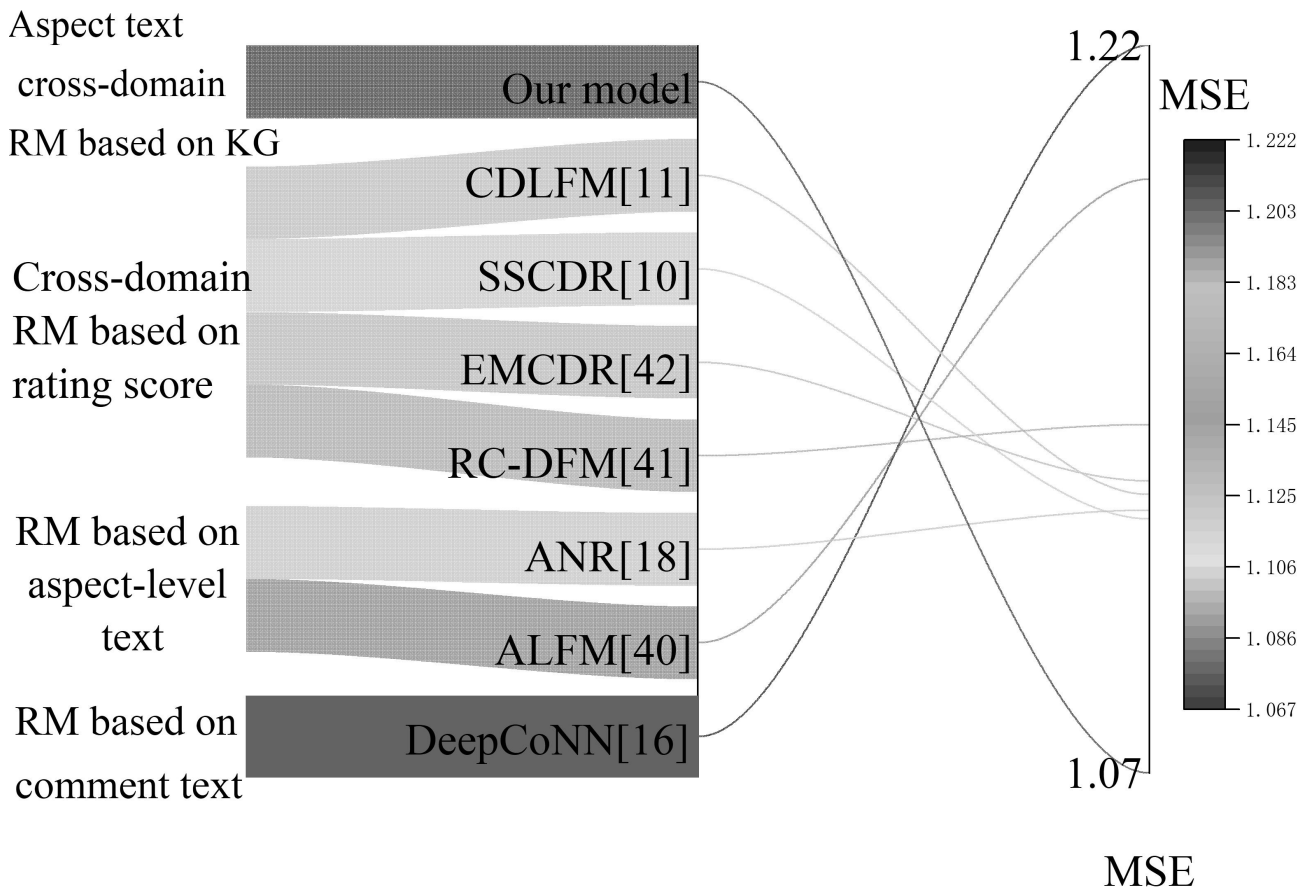


Fig. 9. The MSE comparison on the Book-Music dataset

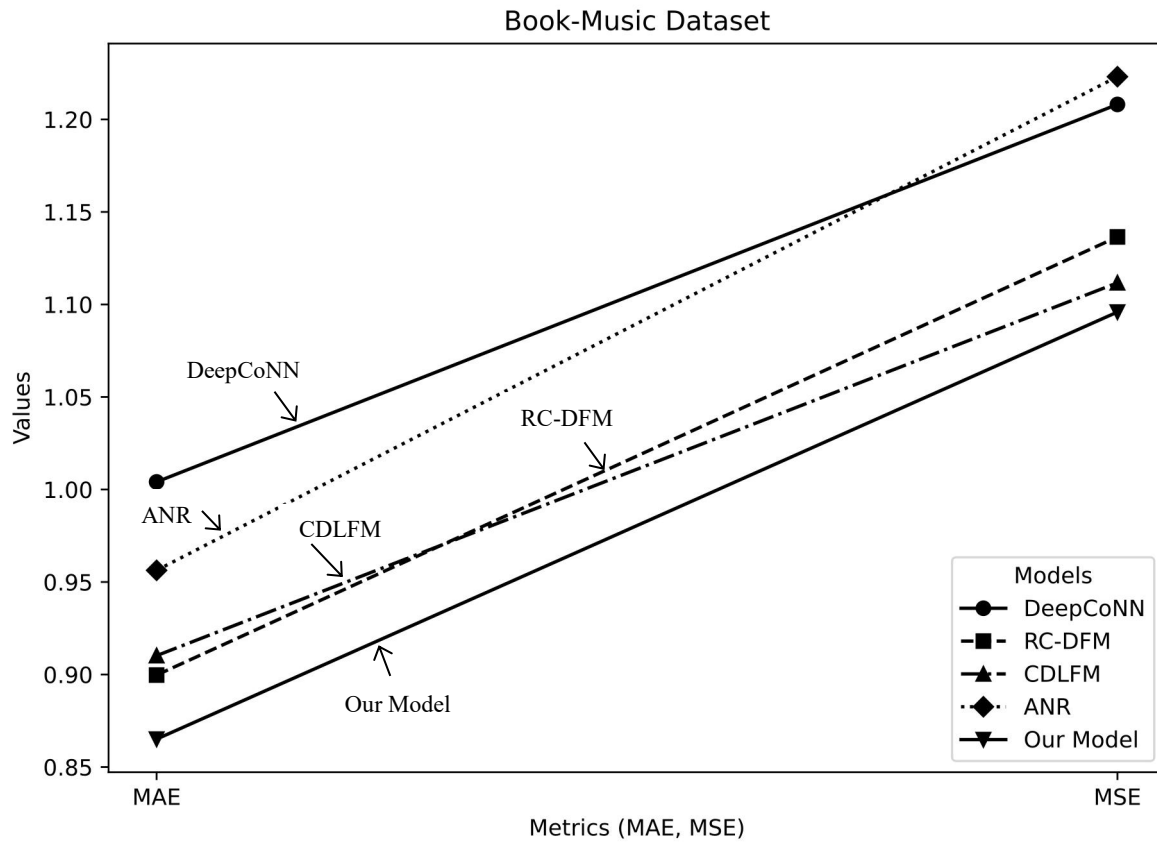


Fig. 10. Performance Comparison on the Book-Music Dataset

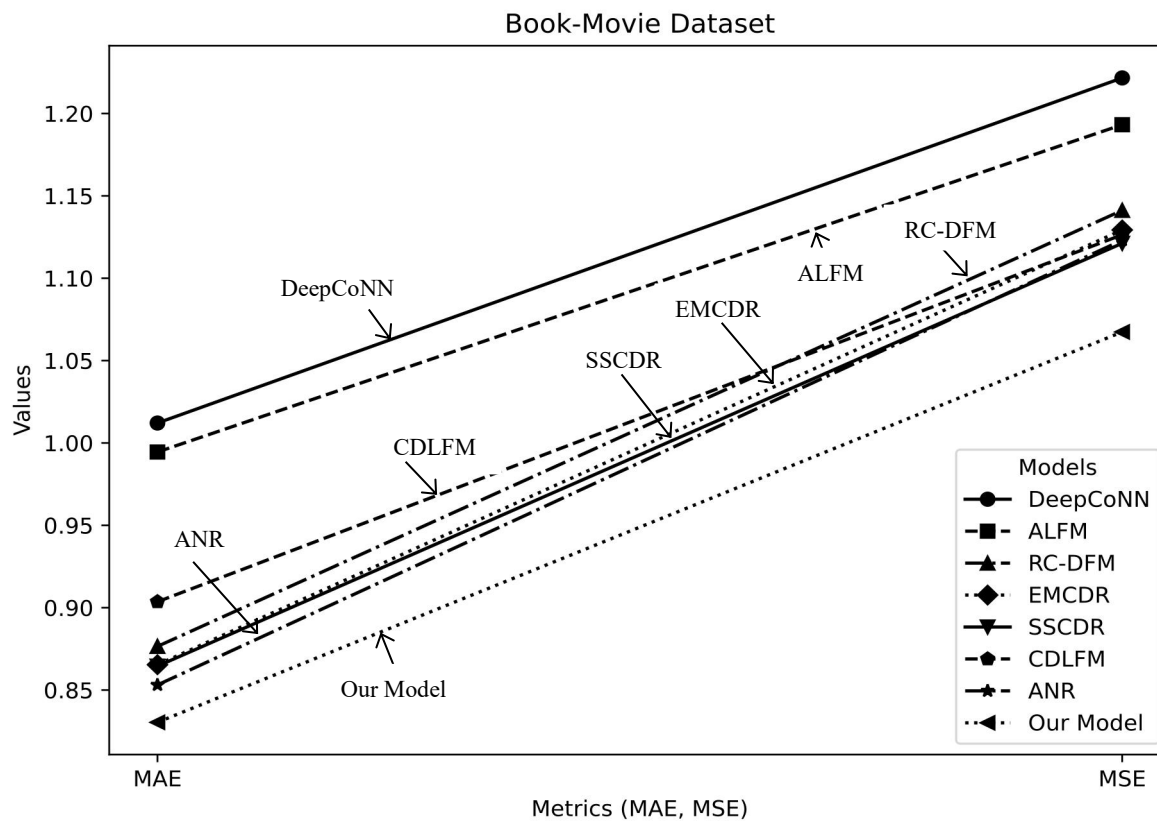


Fig. 11: Performance Comparison on the Book-Movie Dataset

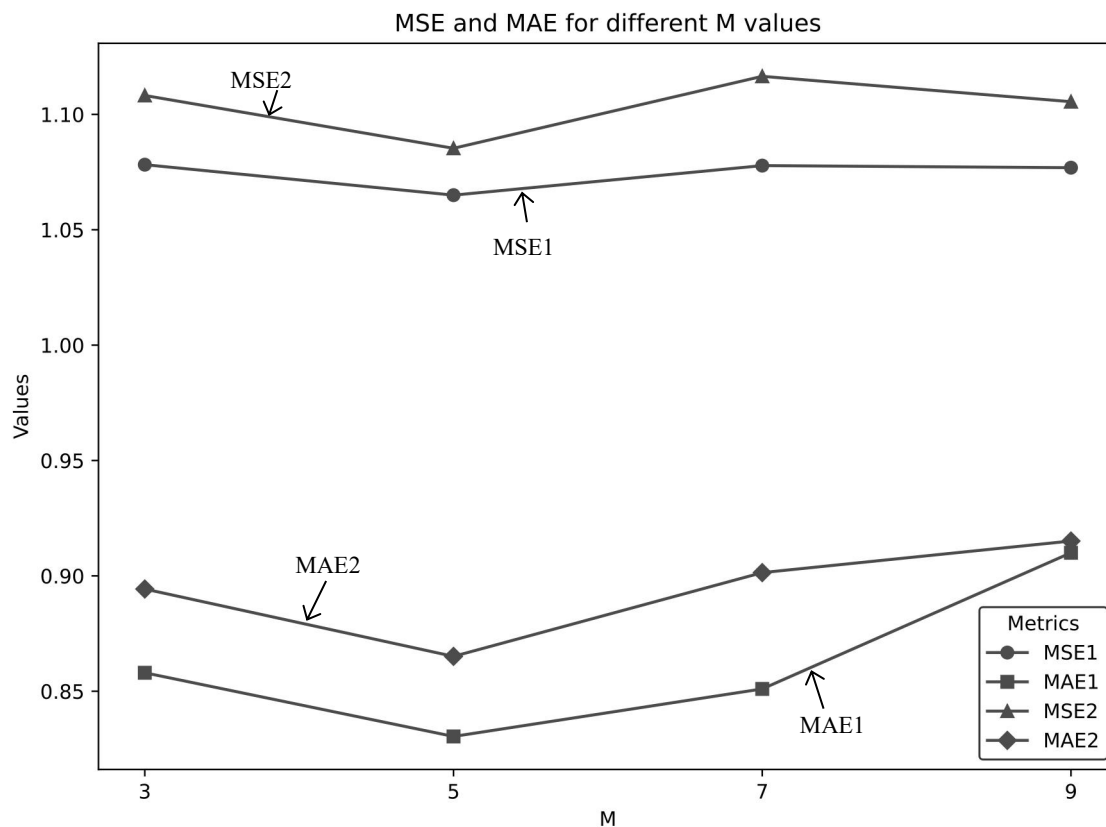


Fig. 12. We set the number of aspect quantity M to 3,5,7,9. On the book-movie and book-music data set, we experimentally analyzed the influence of aspect quantity M on the model performance.

As shown in Table II, our model's MAE and MSE were significantly lower than those of DeepCoNN[10], with performance improvements of 18% and 12.62%, respectively. DeepCoNN directly converted user and item documents into vectors when modeling user preferences, focusing only on single feature representation and failing to capture diversified user and item features. However, not all comment sections are equally important, and the same words may have different meanings. ANR and our model delved deeply into the features, extracting aspect features and verifying the role of aspect-level text analysis.

Our model's performance on the book-movie dataset was significantly better than that of EMCDD[29], SSCDR[5] and CDLFM[6], with MSE performance improving by 5.48%, 4.8%, and 5.24%, respectively. SSCDR outperformed EMCDD, because the user-assisted information in SSCDR enriched the cross-domain interactive data. RC-DFM[28] extended the stacked denoising autoencoder to effectively integrate comment and item content with the scoring matrix of auxiliary and target domains. However, it only captured user and item features in a broad sense and did not deeply explore the semantic relationships between them. KG can capture rich semantic information between users and items, making it useful for obtaining user auxiliary information. On the book-movie dataset, compared with RC-DFM, the MAE and MSE performance of our model improved by 5.26% and 6.47%, respectively. On the book-music dataset, the performance improved by 3.85% and 3.57%, respectively. The result indicate that using KG to obtain user auxiliary information can effectively improve the performance of the cross-domain model. The performance comparison on the

Book-Movie dataset between different model types is shown in fig.8 and fig.9.

Fig. 10 and Fig. 11 compare the performance of various recommendation models on the Book-Music and Book-Movie datasets, respectively, using MAE and MSE as evaluation metrics. In both figures, Our Model consistently outperforms all other models, achieving the lowest MAE and MSE values. In Fig. 10 (Book-Music), Our Model shows superior performance compared to models like DeepCoNN and ANR, while in Fig. 11 (Book-Movie), it continues to lead in both MAE and MSE. Overall, these results highlight the superior accuracy and robustness of Our Model, which demonstrates better prediction accuracy and lower error values across both datasets compared to other state-of-the-art models.

The experimental results demonstrate that our model outperforms others, confirming that the extraction of aspect features from comments and the acquisition of user auxiliary information through KG can enrich cross-domain interactive data and enhance recommendation performance.. Compared with the comparison models, our model shows significant improvements on both datasets, which also demonstrates its strong generalization performance.

D. Influence of Hyper Parameters

This part experimentally evaluates the influence of aspects quantity M . According to theory, the larger of M , the finer the division of explanatory aspects, which should improve the model. However, not all aspects are closely related to the field. Therefore, M will only affect the number of aspects in the source and target domain, with minimal impact on the

transfer of user preferences. As shown in Fig. 12, on the book-movie and book-music datasets, the fluctuation of MSE and MAE was small with different M, with the lowest values occurring when M=5. This indicates that aspect granularity influences the model's performance. When the aspect quantity is too small, the feature granularity is too large, making it difficult to accurately capture user and item features. When the number of aspects is too large, the granularity becomes too fine, and not every aspect is closely related to the domain, which can negatively affect the model.

On the two datasets, the line graphs generated from the experimental results show minimal overall variation and tend to be horizontal. The experimental results demonstrate that the number of aspects.

V. CONCLUSION

The cross-domain recommendation model (RM) based on KG alleviates the cold start and data sparsity issues, enabling cross-domain recommendations from books to movies and music. Firstly, aspect-level emotion analysis was used to extract the features of users and items. Then, a book KG was constructed to fully capture similar users and visually display them. Finally, comments from similar users, extracted via KG, were used as auxiliary information and integrated into the cross-domain RM to enrich the interactive data. The experiment has demonstrated that the model's recommendation performance was significantly improved on the Amazon dataset.

Although the model alleviates the data sparsity problem to some extent, there is still room for further improvement. In the future, a KG of book comment text can be constructed, and information from the comment text can be directly mined using KG, rather than relying on structured data. The model's application scenario currently focuses on single-target cross-domain recommendations, specifically from books to movies and music. In the future, the model can be enhanced to apply more broadly to two-goal and multi-goal cross-domain recommendations. As the number of cross-domains increases, the negative transfer effect becomes inevitable, and further in-depth research can be conducted.

REFERENCES

- [1] Yongpeng Wang, "Research on Cross-Domain Recommendation Based on Overlapping and Non-overlapping Users," ENG.D. dissertation, Chongqing University of Posts and Telecommunications, Chongqing, Sichuan, China, 2021.
- [2] Yu Wang, "Cross-domain Recommendation Model Based on User Reviews," ENG.M. dissertation, Guizhou University, Guiyang, Guizhou, China, 2022.
- [3] Feng Zhu, Yan Wang, Chaochao Chen, Jun Zhou, Longfei Li, and Guanfeng Liu, "Cross-Domain Recommendation: Challenges, Progress, and Prospects," 2021.
- [4] Yehui Zhao, Lin Liu, Hailong Wang, Haiyan Han, and Dongmei Pei, "Research on Knowledge Graph Recommendation System," *Exploration of Computer Science and Technology*, vol.17, no.4, pp771-791, 2023.
- [5] Feng Zhu, Yan Wang, Chaochao Chen, Guanfeng Liu, Mehmet Orgun, and Jia Wu, "A Deep Framework for Cross-domain and Cross-system Recommendations," in the 27th International Joint Conference on Artificial Intelligence (IJCAI '18), Stockholm, Sweden, 2018, pp3711-3717. The paper is available at DOI: 10.48550/arXiv.2009.06215.
- [6] Xinghua Wang, Zhaohui Peng, Senzhang Wang, Philip S. Yu, Wenjing Fu, and Xiaoguang Hong, "Cross-Domain Recommendation for Cold-Start Users via Neighborhood Based Feature Mapping," in International Conference on Database Systems for Advanced Applications (DASFAA 2018), Springer, Cham, 2018, pp158-165.
- [7] Zhao Cheng, Chenliang Li, Rong Xiao, Hongbo Deng and Aixin Sun, "CATN: Cross-Domain Recommendation for Cold-Start Users via Aspect Transfer Network," 2020, The paper is available at <https://arxiv.org/pdf/2005.10549v2>.
- [8] Yongqian Lu, "Research and Application of Cross-domain Recommendation Method Based on Fusion of Auxiliary Information," ENG.M. dissertation, Jiangsu University of Science and Technology, Zhenjiang, Jiangsu, China, 2022.
- [9] Nan Xiao, "Research and Implementation of Cross-domain Recommendation Algorithm Based on Fusion of Auxiliary Information," ENG.M. dissertation, Chongqing University of Posts and Telecommunications, Chongqing, Sichuan, China, 2022.
- [10] Lei Zheng, Vahid Noroozi, and Philip s. Yu, "Joint Deep Modeling of Users and Items Using Reviews for Recommendation," in the Tenth ACM International Conference on Web Search and Data Mining (WSDM 2017), Cambridge, United Kingdom, 2017, pp425-434.
- [11] Sungyong Seo, Jing Huang, Hao Yang, and Yan Liu, "Interpretable Convolutional Neural Networks with Dual Local and Global Attention for Review Rating Prediction," in the Eleventh ACM Conference on Recommender Systems (RecSys '17), New York, NY, USA, 2017, pp297-305.
- [12] Jin Yao Chin, Kaigi Zhao, Shafiq Joty, and Gao Cong, "ANR: Aspect-based Neural Recommender," in the 27th ACM International Conference on Information and Knowledge Management, Torino, Italy, 2018, pp147-156. The paper is available at DOI: 10.1145/3269206.3271810.
- [13] Hongtao Liu, Fangzhao Wu, Wenjun Wang, etc., "NRPA: Neural Recommendation with Personalized Attention," in the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR 2019), Paris, France, 2019, pp1233-1236. The paper is available at DOI: 10.1145/3331184.3331371.
- [14] Tay, Yi, Anh Tuan Luu and Siu Cheung Hui, "Multi-Pointer Co-Attention Networks for Recommendation," in the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining (KDD '18), London, United Kingdom, 2018, pp2309-2318. The paper is available at DOI: 10.1145/3219819.3220086.
- [15] Qiang Cui, Tao Wei, Yafeng Zhang, and Qing Zhang, "HeroGRAPH: A Heterogeneous Graph Framework for Multi-Target Cross-Domain Recommendation," in the 14th ACM Conference on Recommender Systems (RecSys '20), Online, Worldwide, 2020.
- [16] Yuting Ye, Xuwu Wang, Jiangchao Yao, Kunyang Jia, Jingren Zhou, Yanghua Xiao and Hongxia Yang, "Bayes EMBEDDING (BEM): Refining Representation by Integrating Knowledge Graphs and Behavior-specific Networks," in the Proceedings of the 28th ACM International Conference on Information and Knowledge Management (CIKM), Beijing, pp679-688, 2019. The paper is available at DOI: 10.1145/3357384.3358014.
- [17] Yixin Cao, Xiang Wang, Xiangnan He, etc., "Unifying Knowledge Graph Learning and Recommendation: Towards a Better Understanding of User Preferences," in the WWW, New York, 2019, pp151-161. The paper is available at DOI: 10.1145/3308558.3313705.
- [18] Lu Yang, and Mingxiang He, "A Chinese Text Sentiment Analysis Model Based on Gated Mechanism and Convolutional Neural Network," *Journal of Computer Applications*, vol. 41, no.10, pp2842-2848, 2021.
- [19] Wenqiang Xu, "Research on Chinese-English Entity relationship joint extraction method based on attention and gating mechanism," ENG.M. dissertation, Southwest University, Chongqing, Sichuan, China, 2022.
- [20] Yann N. Dauphin, Angela Fan, Michael Auli, and David Grangier, "Language Modeling with Gated Convolutional Networks," *arXiv:1612.08083*, 2016. The paper is available at <https://doi.org/10.48550/arXiv.1612.08083>.
- [21] Yiqiu Fang, Zhenkun Zhang, and Junwei Ge, "Cross-domain Recommendation Algorithm Based on Self-attention Mechanism and Transfer Learning," *Computer Science*, vol.49, no.08, pp70-77, 2022.
- [22] Zhuolan Wang, Yuqi Zhang, Mingyu Chen, Yiqi Su, Kai Song, and Chunling Dong, "Research on Construction of Film Knowledge Graph and Film Recommendation Based on Neo4j Graph Database," *Modern Film Technology*, vol.03, pp29-36, 2022.
- [23] Yizhou Sun, Jiawei Han, Xifeng Yan, Philip S. Yu, and Tianyi Wu, "Pathsim: Meta Path-Based Top-K Similarity Search in Heterogeneous Information Networks," *Proceedings of the VLDB Endowment*, vol.4, no.11, pp992-1003, 2011. The paper is available at DOI: 10.14778/3402707.3402736.
- [24] Qian Zhou, Xun Wang, Linghui Li, Shucheng Huang, and Yunmarsha Wang, "Book Recommendation Algorithm Combining Knowledge

- Graph and Collaborative Filtering,” Journal of Software Guide, vol.21, no.08, pp56-61,2022.
- [25] Guoliang Ji, Shizhu He, Liheng Xu, Kang Liu, and Jun Zhao. “Knowledge Graph Embedding via Dynamic Mapping Matrix,” In Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing, Beijing, China, vol.1, pp687-696,2015. The paper is available at <https://aclanthology.org/P15-1067>
- [26] Libing Wu , Cong Quan, Chenliang Li, Qian Wang, and Bolong Zheng. “A Context-Aware User-Item Representation Learning for Item Recommendation,” ACM Transactions on Information Systems (TOIS), vol. 37, pp1-29, 2017. The paper is available at DOI: 10.48550/arXiv.1712.02342.
- [27] Zhiyong Cheng, Ying Ding, Lei Zhu and M. Kankanhalli. “Aspect-Aware Latent Factor Model: Rating Prediction with Ratings and Reviews,” in the World Wide Web Conference (WWW), Lyon, France, 2018, pp639-648. The paper is available at DOI: 10.1145/3178876.3186145.
- [28] Wenjing Fu, Zhaohui Peng, Senzhang Wang, Yang Xu, and Jin Li, “Deeply Fusing Reviews and Contents for Cold Start Users in Cross-Domain Recommendation Systems,” in the AAAI Conference on Artificial Intelligence (AAAI), Hawaii, United States, vol 33, no.1, pp94-101,2019. The paper is available at DOI: 10.1609/aaai.v33i01.330194.
- [29] Tong Man, Huawei Shen, Xiaolong Jin, and Xueqi Cheng, “Cross-Domain Recommendation: An Embedding and Mapping Approach,” in the Twenty-Sixth International Joint Conference on Artificial Intelligence (IJCAI). Melbourne, Australia, 2017,pp2464-2470. The paper is available at DOI: 10.24963/ijcai.2017/343.