

MKGRec: A Meta-Knowledge Guided Multi-View Recommendation Framework via Integrating LLMs

XiangQin Pang¹, ZhenDong Wu^{1,*}, Yang Yang¹, Hetao Chen¹, Jingzhi Zhang²

¹College of Computer Science, Sichuan Normal University, Chengdu, China

²College of Computer Science and Software Engineering, Shenzhen University, Shenzhen, China
wuzd@sicnu.edu.cn

Abstract—To address the limitations of traditional recommender systems in dynamic interest modeling, multi-source semantic fusion, and semantic preservation during augmentation, this paper proposes a Meta-Knowledge Guided Multi-View Recommendation Framework (MKGRec) via integrating Large Language Models. The core of our approach is to construct a meta-knowledge guided multi-view contrastive learning framework. First, we utilize Large Language Models to parse user behavior sequences and infer latent dynamic interest evolution paths. By filtering noise through mutual information maximization, we construct a dynamic user interest graph. Next, guided by meta-knowledge, we perform cross-view semantic alignment and conflict pruning on the collaborative, interest, and knowledge views to achieve efficient fusion of multi-source heterogeneous information. Finally, we construct a semantic-preserving data augmentation pipeline via a dynamic masking strategy. Experiments on three public datasets demonstrate that our method improves Recall@50 and NDCG@50 by up to 6.87% and 13.16% compared to the best baselines. Ablation studies validate the key contributions of the LLM interest inference, meta-knowledge guidance, and interest augmentation learning modules. Furthermore, noise injection experiments confirm the superior robustness of our method. Our code is available at <https://github.com/pxq1/MKGRec>.

Index Terms—Recommender System, Large Language Model, Meta-knowledge, Knowledge Graph, Graph Contrastive Learning

I. INTRODUCTION

Recommender systems serve as a core technology for information filtering. By mining latent associations between user historical behaviors and item features, they have become essential tools for alleviating information overload and enhancing user experience. Traditional collaborative filtering methods rely on similarities within the user-item interaction matrix to make recommendations. However, their performance is often severely limited by data sparsity and cold-start issues [1]. To address these challenges, researchers have introduced Matrix Factorization, content-based methods, and Graph Neural Networks (GNNs) to learn richer representations of users and items [2]. Specifically, Graph Contrastive Learning (GCL) improves model robustness and representation quality by constructing multiple views through data augmentation [3]. Meanwhile, Large Language Models (LLMs) offer a new approach to building interpretable user interest profiles due to

their powerful semantic generation and reasoning capabilities [4]. Additionally, Knowledge Graphs (KGs) provide rich item attributes and relational information to strengthen the semantic foundation of recommendations [5].

Although prior studies have attempted to integrate these technologies, existing methods still face significant limitations in achieving deep, robust, and dynamic semantic fusion. These limitations are primarily observed in three aspects. First, dynamic interest modeling is insufficient. Static embedding methods struggle to capture the temporal evolution of user interests [6]. While LLMs offer rich semantics, they are often susceptible to hallucinations and noise interference [7]. Second, the fusion of heterogeneous multi-source semantics is inefficient. The complex relationships among the implicit semantics of user behavior sequences, the explicit semantics of knowledge graphs, and collaborative signals are not unified in a single model. Existing approaches mostly rely on shallow fusion or manual designs, resulting in semantic gaps and poor interpretability [8]. Third, augmentation strategies often lead to semantic distortion. Common data augmentation techniques in GCL, such as random dropout or masking, easily disrupt the semantic integrity of the graph structure. This introduces misleading signals, causing biases in representation learning and ultimately impairing the accuracy and robustness of the recommender system [9].

To address the challenges mentioned above, this paper aims to explore a core research question: How can we mine deep associations among cross-source data to construct a collaborative guided recommendation framework that integrates LLMs? The objective is to simultaneously achieve accurate dynamic interest modeling, efficient multi-source semantic fusion, and robust augmentation with semantic preservation. To answer this, we propose a Meta-Knowledge Guided Multi-View Recommendation method via integrating LLMs, named MKGRec. The core idea of this paper is to utilize meta-knowledge as a global semantic consensus to coordinate three heterogeneous views: fine-grained interests from LLMs, item semantics from Knowledge Graphs, and collaborative behaviors. By guiding cross-view semantic alignment and pruning conflicts, our method achieves efficient complementarity and deep fusion of multi-source heterogeneous information within contrastive learning.

The contributions of this paper are summarized as follows:

*Corresponding author.

- Constructing an LLM-driven dynamic interest graph: We utilize LLMs to parse user behavior sequences and generate interest texts. By filtering noise via mutual information, we construct a high-quality user dynamic interest graph, thereby enhancing the dynamic embedding representation capability of the model.
- We propose a meta-knowledge guided multi-view contrastive framework. This framework unifies the topological structures of collaborative, interest, and knowledge views using meta-knowledge. By employing a meta-view to guide cross-view semantic alignment and prune conflicting edges, we achieve efficient fusion of multi-source heterogeneous information. Additionally, we introduce dynamic masking and residual reconstruction mechanisms to construct a semantic-preserving data augmentation pipeline, effectively reducing semantic distortion.
- Establishing a systematic experimental validation framework: Extensive experiments on three public datasets demonstrate that MKGRec significantly outperforms mainstream baselines in terms of recommendation accuracy and robustness. Specifically, the average NDCG@50 exceeds baselines by 13.16%. Furthermore, we deeply analyze the model mechanism through ablation studies and parameter sensitivity research.

II. RELATED WORK

A. LLM and KG Integrated Recommendation Methods

Integrating Large Language Models (LLMs) with Knowledge Graphs (KGs) represents a cutting-edge paradigm in recommender systems [10]. Existing research primarily follows two paths. The first is Knowledge Graph-enhanced methods, such as KGAT [11] and KGIN [12]. These methods enrich user and item representations by modeling high-order relational paths, but they rely on static knowledge, making it difficult to capture the dynamic evolution of user interests and fine-grained preferences [13]. The second is LLM-driven methods, which leverage the powerful semantic understanding capabilities of LLMs or collaborate with KGs for retrieval-augmented reasoning to enhance recommendation performance, such as K-RagRec [14] and CoLaKG [10]. Nevertheless, existing fusion methods still face critical challenges. A survey [15] points out that LLMs are prone to introducing noise and hallucinations when utilizing knowledge, whereas most current fusion methods lack systematic noise filtering and reliability control mechanisms [16]. Furthermore, existing works fail to fully bridge the "semantic gap" between textual interests generated by LLMs and structured knowledge in KGs. These limitations constitute the core issues this study aims to address.

B. GCL-based Recommendation Methods

Graph Contrastive Learning (GCL) has become a mainstream paradigm in the recommendation field due to its ability to alleviate data sparsity and cold-start problems through self-supervised signals [17]. Existing works are primarily categorized into general GCL and multi-view GCL [18]. General GCL methods typically construct augmented views based on

a single original graph [19]. For instance, models like SGL [20] and SimGCL [21] generate views mainly through random structural perturbations or noise injection. Although these methods improve robustness, such "semantics-agnostic" strategies easily disrupt key interaction paths and cause semantic drift. In contrast, multi-view GCL constructs multiple semantic graphs by incorporating external knowledge. Studies such as KGCL [22] and MCCLK [23] introduce external knowledge or multi-view collaboration to preserve semantic integrity during the contrastive process. While these methods effectively capture cross-domain information, optimizing them in a pairwise manner may overlook the meta-knowledge regarding global correlations among all views, which limits the deep fusion of heterogeneous information [24].

III. METHOD

A. Problem Definition and Model Overview

Let $\mathcal{U} = \{u_1, u_2, \dots, u_m\}$ be the set of users and $\mathcal{V} = \{v_1, v_2, \dots, v_n\}$ be the set of items. We construct a user-item interaction graph $\mathcal{G}_{cf} = (\mathcal{V}_{cf}, \mathcal{E}^+)$, where $\mathcal{V}_{cf} = \mathcal{U} \cup \mathcal{V}$ and $\mathcal{E}^+$ represents the set of all positive feedback interactions. Additionally, we construct a knowledge graph $\mathcal{G}_K = \{(h, r, t) | h, t \in \mathcal{V}_E, r \in \mathcal{R}\}$, where $\mathcal{V}_E$ denotes the set of entities and $\mathcal{R}$ denotes the set of relations. The objective is to learn a scoring function $f_{rec}(u, v)$ to predict the preference score $y_{u,v}$ of user u for item v , which is then used to rank items for recommendation.

The MKGRec framework consists of four core modules: LLM dynamic interest graph construction, meta-knowledge guided multi-view contrastive learning, user interest augmentation, and a joint optimization objective, as illustrated in Fig. 1.

B. LLM Dynamic Interest Graph Construction

To address the limitations of dynamic interest modeling, this section proposes the LLM Dynamic Interest Graph Construction. We utilize LLMs to infer fine-grained temporal user interests and filter semantic noise based on the Mutual Information Maximization (MIM) criterion. The aim is to construct a dynamic user interest graph with accurate representations.

1) *Dynamic Interest Graph Construction*: To address the inability of traditional static representations to capture the temporal evolution of interests, we leverage the semantic understanding and reasoning capabilities of LLMs. Specifically, we infer fine-grained and scalable dynamic interest representations from user historical interaction sequences.

First, we define the historical interaction sequence of user u , sorted by real interaction timestamps, as $V_u = [v_1, v_2, \dots, v_n]$, where v_i represents the item from interaction i . We introduce a text mapping function $f_{meta}(\cdot)$ to transform items into natural language descriptions using entity attributes from the

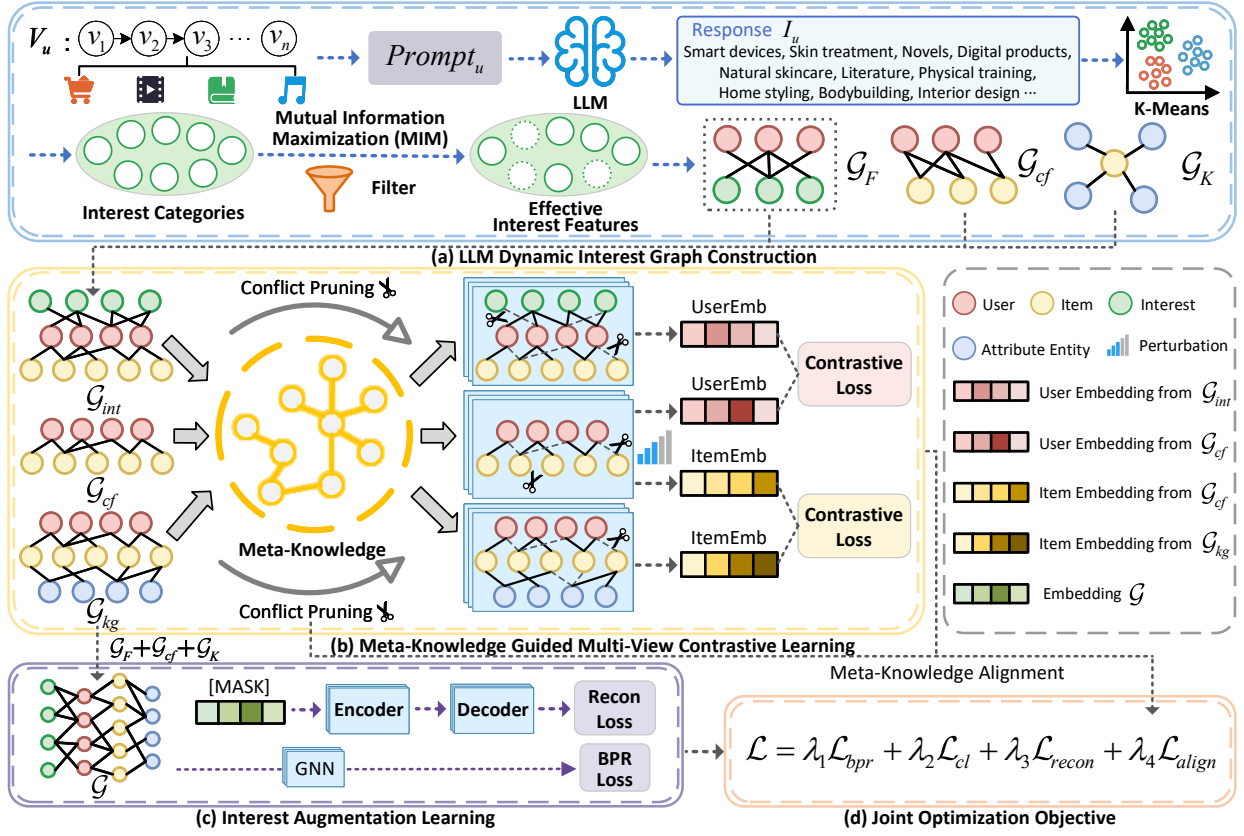


Fig. 1. The overview of the proposed model MKGRec.

knowledge graph. We then construct the prompt text $Prompt_u$ as follows:

$$\begin{aligned}
 Prompt_u = & \text{"User interaction sequence : ["} \\
 & + \text{Concat}(f_{meta}(V_u)) + \text{"}\backslash n\text{"} + \\
 & \text{"Task : Infer fine-grained dynamic interests."}
 \end{aligned} \quad (1)$$

We then utilize the LLM to generate user interest text:

$$I_u = \text{LLM}(Prompt_u), \quad (2)$$

To address potential semantic redundancy in the generated text, we construct a global interest set $\mathcal{I} = \bigcup_u I_u$ and employ the K-Means clustering algorithm $f_{clust}(\cdot)$ to merge similar descriptions. We then extract K representative interest category centers, denoted as L :

$$L = \{l_1, l_2, \dots, l_M\} = f_{clust}(\mathcal{I}, K), \quad (3)$$

where K denotes the number of clusters.

2) *Interest Feature Selection via MIM*: Given that LLM-generated texts often contain hallucination noise, we introduce the MIM criterion to select highly discriminative interest features. By constructing the interest graph using items with high mutual information, we strengthen the dependency between features and categories. This approach significantly enhances the effectiveness and robustness of the representations.

Definition III.1. Let random variable X be the set of candidate interest features, and L be the set of clustered interest categories. For any feature $x \in X$ and category $l \in L$, the Pointwise Mutual Information between them is defined as:

$$PMI(x, l) = \log \frac{p(x, l)}{p(x)p(l)}, \quad (4)$$

where $p(x, l)$ denotes the joint probability, and $p(x)$ and $p(l)$ are the marginal probabilities, respectively.

To filter out generic words with low mutual information or weakly correlated noise, we define a feature selection function. This function retains only the feature items whose correlation with any interest category exceeds the mutual information threshold θ . The filtered set of effective interest features $\mathcal{F}$ is denoted as:

$$\mathcal{F} = \{x \in X \mid \max_{l \in L} PMI(x, l) > \theta\}, \quad (5)$$

Finally, we associate the filtered interest feature set $\mathcal{F}$ with the user nodes $\mathcal{U}$ to construct the user interest graph $\mathcal{G}_F$:

$$\mathcal{G}_F = (\mathcal{V}_F, \mathcal{E}_F), \quad \mathcal{V}_F = \mathcal{U} \cup \mathcal{F}, \quad (6)$$

Where, if user u possesses the interest feature $f \in \mathcal{F}$, an edge $(u, f) \in \mathcal{E}_F$ is formed.

Both LLM inference and clustering are performed offline. The complexity of the online graph construction is $O(|\mathcal{U}| \cdot |\mathcal{F}|)$. Algorithm 1 presents the pseudocode.

Algorithm 1 LLM Dynamic Interest Graph Construction

Input: User set $\mathcal{U}$ and interaction history $\{V_u\}_{u \in \mathcal{U}}$; LLM Model $\text{LLM}(\cdot)$; Clustering function $f_{\text{clust}}(\cdot)$; Cluster number K ; PMI Threshold θ .

Output: User Interest Graph $\mathcal{G}_F$.

```
1: 1. Dynamic Interest Reasoning:
2: for each  $u \in \mathcal{U}$  do
3:   Construct prompt text  $\text{Prompt}_u$  and generate interest
   description
4:    $I_u \leftarrow \text{LLM}(\text{Prompt}_u)$ 
5: end for
6: 2. Global Interest Clustering:
7:  $\mathcal{I} \leftarrow \bigcup_{u \in \mathcal{U}} I_u$ 
8:  $L \leftarrow f_{\text{clust}}(\mathcal{I}, K)$ 
9: 3. Feature Selection based on PMI Maximization:
10: Extract candidate feature set  $X$  from  $\mathcal{I}$ 
11:  $\mathcal{F} \leftarrow \{x \in X \mid \max_{l \in L} \text{PMI}(x, l) > \theta\}$ 
12: 4. Construct User Interest Graph:
13: Initialize  $\mathcal{G}_F = (\mathcal{U} \cup \mathcal{F}, \mathcal{E}_F)$ 
14: for each  $u \in \mathcal{U}, f \in \mathcal{F}$  do
15:   if  $u$  possesses  $f$  then  $(u, f) \in \mathcal{E}_F$ 
16: end for
17: return  $\mathcal{G}_F$ 
```

C. Meta-Knowledge Guided Multi-View Contrastive Learning

To address the challenges of multi-source semantic heterogeneity and noise interference, this section proposes a meta-knowledge guided multi-view contrastive learning mechanism. We define "meta-knowledge" as a global semantic consensus extracted from multi-source heterogeneous views. Its purpose is to bridge the semantic gap between different views. By guiding cross-view semantic alignment and conflict resolution, it achieves the robust fusion of multi-source information.

1) *Heterogeneous View Construction:* To overcome the semantic sparsity of a single view, we construct three heterogeneous views to achieve multi-dimensional information complementarity. First, we retain the basic user-item interaction graph as the Collaborative View to preserve core interaction signals. Second, we fuse the user interest graph with the interaction graph to build the Interest View, which captures fine-grained evolution. Finally, we integrate the external knowledge graph into the interaction graph to form the Knowledge View, thereby introducing entity attribute correlations. The three heterogeneous views are defined as follows:

$$\begin{aligned} \mathcal{G}_{cf} &= (\mathcal{V}_{cf}, \mathcal{E}_{cf}), \\ \mathcal{G}_{int} &= (\mathcal{V}_{int}, \mathcal{E}_{int}), \quad \mathcal{V}_{int} = \mathcal{U} \cup \mathcal{V} \cup \mathcal{F}, \quad \mathcal{E}_{int} = \mathcal{E}_{cf} \cup \mathcal{E}_F, \\ \mathcal{G}_{kg} &= (\mathcal{V}_{kg}, \mathcal{E}_{kg}), \quad \mathcal{V}_{kg} = \mathcal{U} \cup \mathcal{V} \cup \mathcal{V}_E, \quad \mathcal{E}_{kg} = \mathcal{E}_{cf} \cup \mathcal{E}_K, \\ z^{cf} &= f_G^l(Z, \mathcal{G}_{cf}), \quad z^{int} = f_G^l(Z, \mathcal{G}_{int}), \quad z^{kg} = f_G^l(Z, \mathcal{G}_{kg}), \end{aligned} \quad (7)$$

where $f_G^l(\cdot)$ denotes an l -layer Graph Neural Network encoder, and Z represents the initial node embeddings.

2) *Meta-knowledge Extraction:* To bridge the semantic discrepancy among heterogeneous views, this module constructs a unified meta view. It achieves this by exploring the global semantic consensus across different views. This process

ultimately facilitates knowledge transfer. First, we build the meta-view based on the Interest View $\mathcal{G}_{int}$, Collaborative View $\mathcal{G}_{cf}$, and Knowledge View $\mathcal{G}_{kg}$, utilizing it as a unified meta-representation for cross-view alignment. Second, we design a meta-knowledge learning network for each view, composed of two fully connected layers with ReLU activation functions. Let $H^{(k)}$ denote the view matrix. The transformation matrix is calculated as follows:

$$P_{G_1}^{(k)} = \text{MLP}_1^{(k)}(H^{(k)}), \quad P_{G_2}^{(k)} = \text{MLP}_2^{(k)}(H^{(k)}), \quad (8)$$

where MLP_1 evaluates the importance of the global structure, and MLP_2 projects the representations into the meta semantic space. The transformation matrix is dynamically generated based on the node features within the view, enabling the autonomous learning of meta-knowledge.

Finally, we construct the meta-view representation by applying an attention aggregation operation at the view level:

$$\mathbf{H}_{\mathcal{M}} = \sum_{k=1}^K \sigma(g^{(k)} P_{G_1}^{(k)}) \cdot P_{G_2}^{(k)}, \quad K=3 \quad (9)$$

where the weight $g^{(k)}$ reflects the contribution of each view to the meta-view, and $\sigma(\cdot)$ denotes the activation function.

3) *Cross-View Semantic Alignment:* The meta-view representation is constructed as a global semantic coordinate system. It provides a unified metric basis for multi-source heterogeneous views to constrain cross-view semantic alignment. Node correlations within a specific view no longer rely solely on local topology but are recalibrated through meta-knowledge. Specifically, for any node pair n_i and n_j in view k , the model maps them to the meta-semantic space. It utilizes a view weight matrix to modulate the meta-representations and calculates the weighted cosine similarity between them:

$$\mathcal{S}_{ij}^k = \cos(W^{(k)} \odot h_{\mathcal{M},i}, W^{(k)} \odot h_{\mathcal{M},j}) \quad (10)$$

where $\odot$ denotes the Hadamard product, and $h_{\mathcal{M},i}$ and $h_{\mathcal{M},j}$ represent the meta-view representations of nodes n_i and n_j , respectively. $W^{(k)}$ is the trainable weight matrix corresponding to view k , which can be dynamically adjusted based on local features and meta-information.

To ensure that the edge weights are non-negative, the model normalizes the similarity into the edge confidence r_{ij}^k :

$$r_{ij}^k = (\mathcal{S}_{ij}^k + 1)/2 \quad (11)$$

We filter out edges with weights lower than the pruning threshold ϵ from the view:

$$\bar{A}_{ij}^k = \begin{cases} 1, & r_{ij}^k \geq \epsilon \text{ and } A_{ij}^k = 1, \\ 0, & \text{else,} \end{cases} \quad k=3 \quad (12)$$

Here, ϵ is a non-negative threshold adaptively set based on the statistics of meta-view similarity to prune low-confidence connections. A_{ij}^k denotes the adjacency relationship between node n_i and n_j in the k -th view. Guided by the structural prior weights provided by meta-knowledge, this filtering process identifies and removes conflicting edges and redundant associations, achieving cross-view semantic alignment.

4) *Multi-view Contrastive Learning*: To address the semantic discrepancy between views and the sparsity of collaborative signals, we propose a cross-view contrastive learning mechanism using the Collaborative View as a pivot. This method utilizes self-supervised signals to achieve implicit semantic alignment. Specifically, we enforce representation consistency for the same user across the Collaborative View $\mathcal{G}_{cf}$ and the Interest View $\mathcal{G}_{int}$. Similarly, we constrain the consistency of the same item across the Collaborative View $\mathcal{G}_{cf}$ and the Knowledge View $\mathcal{G}_{kg}$. We construct negative samples through random sampling within the batch.

To align entity representations across different views and enhance their discriminative capability, we introduce an InfoNCE-based contrastive loss:

$$\mathcal{L}_{cl} = \sum_{u \in \mathcal{U}} -\log \frac{\exp(\text{sim}(z_u^{cf}, z_u^{int})/\tau)}{\sum_{u' \in \mathcal{U}} \exp(\text{sim}(z_u^{cf}, z_{u'}^{int})/\tau)} + \sum_{v \in \mathcal{V}} -\log \frac{\exp(\text{sim}(z_v^{cf}, z_v^{kg})/\tau)}{\sum_{v' \in \mathcal{V}} \exp(\text{sim}(z_v^{cf}, z_{v'}^{kg})/\tau)} \quad (13)$$

where $\text{sim}(\cdot, \cdot)$ denotes the cosine similarity, and τ represents the temperature parameter.

5) *Meta-knowledge Alignment Loss*: Although structural pruning optimizes the graph topology, the views still lack explicit consistency constraints in the feature space. To this end, we design a self-distillation meta-knowledge alignment mechanism. We use the meta-view containing global semantics as the "teacher" to guide the feature calibration of the three sub-views (Collaborative, Interest, and Knowledge), which act as "students". By minimizing the representation difference between the sub-views and the global meta-view, we enhance the semantic robustness of each view. The alignment loss function $\mathcal{L}_{align}$ is defined as follows:

$$\mathcal{L}_{align} = \sum_{g \in \mathcal{G}} \left(\sum_{u \in \mathcal{U}} \|e_u^g - \text{sg}(h_{\mathcal{M}, u})\|_2^2 + \sum_{v \in \mathcal{V}} \|e_v^g - \text{sg}(h_{\mathcal{M}, v})\|_2^2 \right) \quad (14)$$

Here, $\mathcal{G} = \{\mathcal{G}_{cf}, \mathcal{G}_{int}, \mathcal{G}_{kg}\}$ denotes the set of heterogeneous views, while e and h represent the node representations of the sub-views and the meta-view, respectively. We introduce a stop-gradient operator $\text{sg}(\cdot)$ to freeze the gradient updates of the meta-view. This forces the sub-views to align with the global semantic distribution. Furthermore, it effectively prevents semantic collapse caused by the meta-view adapting to the noise within the sub-views, thereby ensuring the stability of the unidirectional guidance provided by meta-knowledge.

D. Interest Augmentation Learning

Existing GCL methods often rely on random structural perturbations, which tend to disrupt critical semantic paths and cause structural semantic distortion. To address this, we propose an interest-aware semantic preservation augmentation strategy. Unlike random destruction, we specifically perform a "mask-and-reconstruct" graph auto-encoding task on the

interest nodes. Furthermore, we implement a dynamic masking rate scheduler to achieve stable training in an easy-to-hard manner. To reconstruct interests on a unified message passing backbone, we aggregate collaborative, knowledge, and interest relations into a basic heterogeneous graph:

$$\mathcal{G} = \{(h, r, t) \mid h, t \in \mathcal{N} = (\mathcal{U} \cup \mathcal{V} \cup \mathcal{F}), r \in \mathcal{R}'\} \quad (15)$$

Here, $\mathcal{N}$ encompasses the user, item, and interest nodes, while $\mathcal{R}'$ represents the union of relations across the three views. The initial node embedding matrix is defined as $Z \in \mathbb{R}^{|\mathcal{N}| \times D}$.

To enhance training stability, we introduce a curriculum-based masking rate scheduler. Let the masking rate at the e -th training epoch be:

$$\rho^e = \text{clip}(\Psi(e), \omega), \quad \Psi(e) = \begin{cases} \Psi_{\text{lin}}(e) = \alpha + e \frac{(\omega - \alpha)}{E}, \\ \Psi_{\text{exp}}(e) = \alpha \left(\frac{\omega}{\alpha}\right)^{\frac{e}{E}}, \end{cases} \quad (16)$$

Here, α , ω , and E represent the initial masking rate, the maximum masking rate, and the number of epochs required to reach the maximum masking rate, respectively. $\text{clip}(\cdot)$ denotes the truncation function. When $0 < e < E$, the inequality $\Psi_{\text{exp}}(e) < \Psi_{\text{lin}}(e)$ holds. This indicates that the exponential scheduler exhibits a smoother growth curve, which aligns well with the "easy-to-hard" training paradigm.

Based on the masking rate ρ^e , we randomly sample a subset $\tilde{\mathcal{F}} \subset \mathcal{F}$ and replace the embeddings of the sampled nodes with a learnable mask token $[M]$:

$$z'_f = \begin{cases} z_{[M]}, & f \in \tilde{\mathcal{F}}, \\ z_f, & f \notin \tilde{\mathcal{F}}, \end{cases} \quad (17)$$

The masked embeddings Z' and the base heterogeneous graph $\mathcal{G}$ are fed into the graph encoder $f_{\text{enc}}(\cdot)$ to obtain the intermediate representation Z'' . To alleviate underfitting and stabilize the reconstruction, we introduce a residual feed-forward transformation $\chi(\cdot)$ between the encoder and the decoder to establish an identity residual connection:

$$\begin{aligned} Z'' &= f_{\text{enc}}(Z', \mathcal{G}), \\ \tilde{Z} &= \chi(Z'') + Z'', \quad \chi(Z'') = \text{LN}(\sigma(Z'' W_1) W_2), \\ Z''' &= f_{\text{dec}}(\tilde{Z}, \mathcal{G}), \end{aligned} \quad (18)$$

Here, $\sigma(\cdot)$ denotes the ReLU activation function, and $\text{LN}(\cdot)$ represents Layer Normalization. f_{enc} and f_{dec} refer to the graph encoder and decoder, respectively, which do not share parameters. W_1 and W_2 are the learnable weight matrices.

Finally, we formulate the user interest reconstruction loss function with a focusing parameter η :

$$\mathcal{L}_{\text{recon}} = \frac{1}{|\tilde{\mathcal{F}}|} \sum_{f \in \tilde{\mathcal{F}}} \left(1 - \frac{z_f^T z_f'''}{\|z_f\| \cdot \|z_f'''\|} \right)^\eta, \quad \eta \geq 1, \quad (19)$$

This objective forces the model to reconstruct the masked interest semantics, thereby achieving semantic-preserving augmentation without disrupting the original graph structure.

TABLE I

COMPARISON RESULTS OF MKGRec WITH BASELINE MODELS. **BOLD** AND UNDERLINED VALUES DENOTE THE BEST AND SECOND-BEST RESULTS.

Model	DBbook2014				Book-Crossing				MovieLens-1M			
	R@50	R@100	N@50	N@100	R@50	R@100	N@50	N@100	R@50	R@100	N@50	N@100
LightGCN	0.4214	0.5358	0.2230	0.2481	0.1607	0.2246	0.0765	0.0911	0.4086	0.5574	0.3977	0.4432
SGL	0.4258	0.5249	0.2350	0.2567	0.1613	0.2154	0.0860	0.0982	0.4094	0.5505	0.4004	0.4433
SimGCL	0.4280	0.5369	0.2376	0.2610	0.1653	0.2211	0.0860	0.0988	0.4131	0.5560	0.4048	0.4482
XSimGCL	0.4465	0.5537	0.2441	0.2678	0.1713	0.2290	0.0886	0.1039	0.4223	0.5634	0.4124	0.4512
BIGCF	0.4334	0.5465	0.2383	0.2563	0.1746	0.2335	0.0880	0.1015	0.4245	0.5668	0.4107	0.4473
CKE	0.3419	0.4448	0.1832	0.2060	0.1284	0.1691	0.0708	0.0811	0.3948	0.5421	0.3841	0.4295
CFKG	0.3586	0.4559	0.1627	0.1816	0.1216	0.1771	0.0503	0.0617	0.3930	0.5439	0.3672	0.4142
KGAT	0.4174	0.5309	0.2110	0.2360	0.1564	0.2133	0.0865	0.1004	0.4059	0.5527	0.3928	0.4378
KGIN	0.4147	0.5337	0.2200	0.2462	0.1468	0.2075	0.0681	0.0819	0.4095	0.5581	0.4004	0.4453
KGCL	0.4308	0.5271	0.2453	0.2666	0.1562	0.2049	0.0859	0.0971	0.4083	0.5496	0.3996	0.4430
MCCLK	0.4376	0.5491	0.2054	0.2273	0.1476	0.2039	0.0656	0.0774	0.4176	0.5655	0.3876	0.4339
KGRec	0.4415	0.5497	0.2373	0.2611	0.1473	0.2075	0.0699	0.0836	0.4136	0.5633	0.4072	0.4524
CIKGRc	<u>0.4642</u>	<u>0.5756</u>	<u>0.2494</u>	<u>0.2739</u>	<u>0.1791</u>	<u>0.2455</u>	<u>0.0889</u>	<u>0.1041</u>	<u>0.4250</u>	<u>0.5718</u>	<u>0.4162</u>	<u>0.4609</u>
RLMRec	0.4290	0.5422	0.2291	0.2540	0.1583	0.2218	0.0769	0.0907	0.3809	0.5298	0.3627	0.4123
MKGRec	0.4776	0.5859	0.2633	0.2863	0.1914	0.2557	0.1006	0.1153	0.4305	0.5809	0.4229	0.4706
%IMP.	2.89%	1.79%	5.57%	4.53%	6.87%	4.15%	13.16%	10.76%	1.29%	1.59%	1.61%	2.10%

E. Joint Optimization Objective

The training objective of the model also adopts the Bayesian Personalized Ranking (BPR) loss function $\mathcal{L}_{bpr}$ to optimize the recommendation ranking task:

$$\hat{Z} = f_G^l(Z, \mathcal{G}),$$

$$\mathcal{L}_{bpr} = - \sum_{(u, v, v')} \ln(\sigma(\hat{z}_u^T \hat{z}_v - \hat{z}_u^T \hat{z}_{v'})), \quad (20)$$

Where $\hat{Z}$ represents the final node representation. $(u, v) \in \mathcal{E}^+$ denotes a positive interaction, and v' is a negative sample randomly sampled for user u . $\sigma(\cdot)$ is the sigmoid function.

By integrating the previously defined graph contrastive loss, interest reconstruction loss, and meta-knowledge alignment loss, the overall joint optimization objective is formulated as:

$$\mathcal{L} = \lambda_1 \mathcal{L}_{bpr} + \lambda_2 \mathcal{L}_{cl} + \lambda_3 \mathcal{L}_{recon} + \lambda_4 \mathcal{L}_{align} + \lambda_5 \|\Theta\|_2^2 \quad (21)$$

Here, Θ denotes the set of trainable model parameters, and λ is a hyperparameter that balances the different sub-tasks and regularization weights. We determine the value of λ through grid search on the validation set to balance training convergence speed and overall performance.

IV. EXPERIMENTS

To evaluate the effectiveness of the proposed MKGRec, we conduct extensive experiments to answer the following research questions. RQ1: How does MKGRec perform compared with existing models? RQ2: Are the main components of the model effective? RQ3: Is MKGRec more robust than baseline models against data noise? RQ4: How do different hyperparameter settings affect the performance of MKGRec?

A. Experimental Settings

1) *Datasets and Evaluation Metrics*: We conduct experiments on three public datasets: DBbook2014, Book-Crossing, and MovieLens-1M. These datasets come from different domains and are representative in terms of data scale and sparsity. The statistics of the datasets are shown in Table II. To evaluate

TABLE II
STATISTICS OF EXPERIMENTAL DATASETS.

Dataset	Dbook2014	Book-Crossing	MovieLens-1M
#Users	5,576	6,616	6,040
#Items	2,680	8,853	3,260
#Ratings	65,961	110,662	998,539
#Density	0.44%	0.19%	5.07%
#Relations	13	4	20
#Entities	8,762	1,404	14,377
#Triples	134,223	1,137	415,104

the performance of the MKGRec model in the Top-K item recommendation task, we employ two commonly used metrics: Recall@K (R@K) and NDCG@K (N@K), where K is set to 50 and 100.

2) *Baselines*: To comprehensively validate the performance of MKGRec, we select four categories of mainstream baseline methods for comparison: (i) Classic graph-based recommendation models, including LightGCN [25]. (ii) Self-supervised and contrastive learning models, including SGL [20], SimGCL [21], XSimGCL [26], and BIGCF [27]. (iii) Knowledge graph-enhanced recommendation models, including CKE [28], CFKG [29], KGAT [11], KGIN [12], KGCL [22], MCCLK [23], and KGRec [5]. (iv) Generative and large model-driven recommendation models, including CIKGRc [30] and RLM-Rec [31].

3) *Implementation Details*: Following the baseline settings for sparse datasets [30], we split the datasets into training, validation, and testing sets in a 7:1:2 ratio. The graph encoders all adopt the LightGCN architecture. We employ the Adam optimizer with a batch size of 1024, and the learning rate is searched within the range of $\{5e^{-4}, 8e^{-4}\}$. To prevent overfitting, we apply L_2 regularization to the embeddings of the collaborative view and the knowledge view, with coefficients set to 1×10^{-5} and 1×10^{-3} , respectively. To balance high quality semantic reasoning and the low latency overhead of recommender systems, the LLM component integrates the lightweight DeepSeek R1 (1.5B) model. We set the maximum

number of training epochs to 2000, accompanied by an early stopping mechanism with a patience of 20. The final performance metrics are obtained by averaging the results of 10 independent runs under different random seeds.

B. Main Results (RQ1)

The performance comparison between MKGRec and the baseline models is shown in Table I. MKGRec consistently outperforms all baseline methods across the three recommendation datasets. First, compared to traditional graph models such as LightGCN and KGAT, MKGRec demonstrates significant advantages. This suggests that relying solely on interaction topology mining or static knowledge graphs is difficult to capture users' deep intentions, whereas introducing dynamic interests inferred by LLMs effectively alleviates semantic sparsity. Second, MKGRec achieves a significant improvement of 13.16% in NDCG@50 on the highly sparse Book-Crossing dataset. This gain is notably larger than the improvement on the relatively dense ML-1M dataset. We attribute this to the alignment mechanism for meta-knowledge. In sparse scenarios that heavily depend on external semantics, this mechanism effectively coordinates heterogeneous complementary information. As a result, it successfully overcomes the limitations of shallow fusion. Finally, MKGRec improves NDCG@50 by 5.57% on the DBbook2014 dataset. This validates the effectiveness of the "mask-and-reconstruct" augmentation strategy, which mitigates the semantic distortion problem caused by random data augmentation.

C. Ablation Study (RQ2)

To investigate the contributions of the core components in MKGRec, we construct four variants: w/o IG (removing the interest graph), w/o Meta (removing the meta-knowledge guidance), w/o Int (removing the interest augmentation), and w Random (applying random edge dropping). The ablation results are presented in Table III. The performance degradation of w/o IG confirms that constructing a dynamic interest graph effectively introduces rich semantics. It also captures detailed dynamic preferences and alleviates data sparsity. Furthermore, the R@50 score of w/o Meta on the Book-Crossing dataset drops by 5.85%. This demonstrates orthogonal gains rather than module redundancy. It confirms the critical role of meta-knowledge in bridging the semantic gap among views and suppressing noise. Finally, the performance decline of w/o Int and w Random verifies that the "mask-and-reconstruct" semantic preservation augmentation strategy enhances representation robustness.

D. Robustness Analysis (RQ3)

To validate the robustness, we introduce structural noise with intensity $s \in [0.0, 1.0]$ into the experiments. We then compare our model with the best baseline driven by large language models, CIKGRc. As illustrated in Fig. 2, MKGRec demonstrates superior stability compared to CIKGRc as the noise intensity increases. Specifically, on the Book-Crossing dataset with $s = 0.5$, the Recall@50 of MKGRec drops by

TABLE III
ABLATION STUDY OF MKGREC.

Model	DBbook2014		Book-Crossing		MovieLens-1M	
	R@50	N@50	R@50	N@50	R@50	N@50
w/o IG	0.4759	0.2610	0.1816	0.0933	0.4256	0.4154
w/o Meta	0.4712	0.2581	0.1802	0.0994	0.4226	0.4178
w/o Int	0.4764	0.2615	0.1821	0.0938	0.4231	0.4147
w/o Random	0.4751	0.2623	0.1843	0.0954	0.4247	0.4156
MKGRec	0.4776	0.2633	0.1914	0.1006	0.4305	0.4229

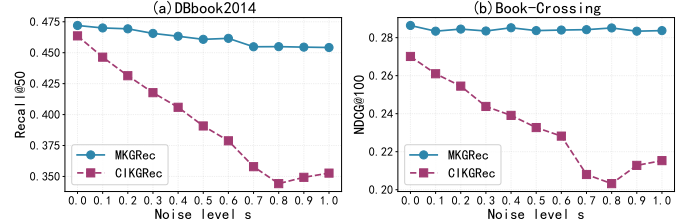


Fig. 2. Performance comparison between MKGRec and CIKGRc under different noise intensities (s).

only 2.37%, which is significantly lower than the 15.70% decline observed in CIKGRc. We attribute this robustness to the mutual information denoising and the conflict pruning and alignment mechanism guided by meta-knowledge. These components effectively resist perturbations in the feature space.

E. Hyperparameter Analysis (RQ4)

1) *Pruning Threshold*: As illustrated in Fig. 3, the model performance initially improves and then declines as the pruning threshold ϵ increases. The model achieves optimal performance within the range of $\epsilon \in [0.3, 0.5]$. This confirms that meta-knowledge serves as a global semantic consensus, enabling the accurate removal of semantic conflicts and redundant edges for efficient cross-view alignment. Conversely, an excessively high threshold ($\epsilon > 0.9$) leads to the erroneous deletion of key topological paths, causing semantic loss. These results indicate that meta-knowledge provides a reliable global metric basis, and an appropriate threshold is crucial for balancing semantic consistency and noise suppression.

2) *Contrastive Loss Temperature*: As shown in Fig. 4, the model achieves optimal performance at $\tau = 0.2$ for DBbook2014 and $\tau = 0.15$ for Book-Crossing. A value of τ that is too small causes the model to focus excessively on hard negative samples, damaging semantic generalization. Conversely, as τ exceeds the threshold, the discriminative power decreases, leading to a decline in performance. A reasonable temperature setting is key to balancing semantic alignment and distribution uniformity.

V. CONCLUSION

This paper proposes MKGRec, a Meta-knowledge Guided Multi-view Recommendation framework integrating Large Language Models. To address the challenges of dynamic interest evolution and heterogeneous information fusion, we

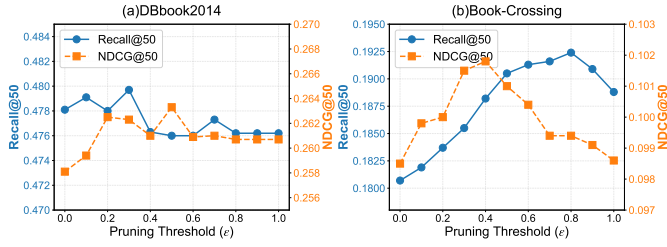


Fig. 3. Impact of Pruning Threshold ϵ .

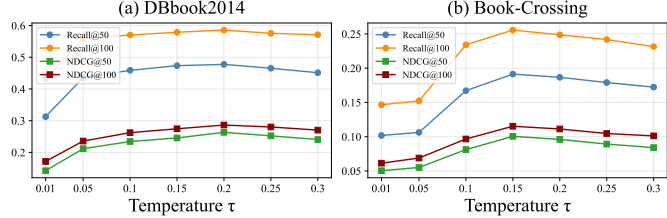


Fig. 4. Impact of Contrastive Learning Temperature τ .

leverage LLMs and mutual information maximization to construct dynamic interest graphs. Furthermore, we achieve cross-view conflict pruning and deep alignment through a meta-knowledge guided mechanism. Meanwhile, we introduce a semantic-preserving augmentation strategy to enhance model robustness. Extensive experiments demonstrate the effectiveness of the MKGRec framework and validate its design. Future work will focus on exploring the transferability and general value of meta-knowledge in cross-domain recommendation.

REFERENCES

- [1] S. Sanner, K. Balog, F. Radlinski, B. Wedin, and L. Dixon, "Large language models are competitive near cold-start recommenders for language-and item-based preferences," in *Proc. ACM RecSys*, 2023, pp. 890–896.
- [2] Y. Wu, L. Zhang, F. Mo, T. Zhu, W. Ma, and J.-Y. Nie, "Unifying graph convolution and contrastive learning in collaborative filtering," in *Proc. ACM SIGKDD*, 2024, pp. 3425–3436.
- [3] Y. Yang, Z. Wu, L. Wu, K. Zhang, R. Hong, Z. Zhang, J. Zhou, and M. Wang, "Generative-contrastive graph learning for recommendation," in *Proc. ACM SIGIR*, 2023, pp. 1117–1126.
- [4] J. Liao, S. Li, Z. Yang, J. Wu, Y. Yuan, X. Wang, and X. He, "Llara: Large language-recommendation assistant," in *Proc. ACM SIGIR*, 2024, pp. 1785–1795.
- [5] Y. Yang, C. Huang, L. Xia, and C. Huang, "Knowledge graph self-supervised rationalization for recommendation," in *Proc. ACM SIGKDD*, 2023, pp. 3046–3056.
- [6] J. Lin, R. Shan, C. Zhu, K. Du, B. Chen, S. Quan, R. Tang, Y. Yu, and W. Zhang, "Rella: Retrieval-enhanced large language models for lifelong sequential behavior comprehension in recommendation," in *Proc. WWW*, 2024, pp. 3497–3508.
- [7] V. Rawte, S. Chakraborty, A. Pathak, A. Sarkar, S. T. I. Tonmoy, A. Chadha, A. Sheth, and A. Das, "The troubling emergence of hallucination in large language models—an extensive definition, quantification, and prescriptive remediations," in *Proc. EMNLP*, 2023, pp. 2541–2573.
- [8] Y. Zhu, L. Wu, Q. Guo, L. Hong, and J. Li, "Collaborative large language model for recommender systems," in *Proc. WWW*, 2024, pp. 3162–3172.
- [9] X. Cai, C. Huang, L. Xia, and X. Ren, "Lightgcl: Simple yet effective graph contrastive learning for recommendation," *arXiv preprint arXiv:2302.08191*, 2023.
- [10] Z. Cui, Y. Weng, X. Tang, F. Lyu, D. Liu, X. He, and C. Ma, "Comprehending knowledge graphs with large language models for recommender systems," in *Proc. ACM SIGIR*, 2025, pp. 1229–1239.
- [11] X. Wang, X. He, Y. Cao, M. Liu, and T.-S. Chua, "Kgat: Knowledge graph attention network for recommendation," in *Proc. ACM SIGKDD*, 2019, pp. 950–958.
- [12] X. Wang, T. Huang, D. Wang, Y. Yuan, Z. Liu, X. He, and T.-S. Chua, "Learning intents behind interactions with knowledge graph for recommendation," in *Proc. WWW*, 2021, pp. 878–887.
- [13] H. Hu, W. Guo, X. Liu, Y. Liu, R. Tang, R. Zhang, and M.-Y. Kan, "User behavior enriched temporal knowledge graphs for sequential recommendation," in *Proc. WSDM*, 2024, pp. 266–275.
- [14] S. Wang, W. Fan, Y. Feng, L. Shanru, X. Ma, S. Wang, and D. Yin, "Knowledge graph retrieval-augmented generation for llm-based recommendation," in *Proc. ACL*, 2025, pp. 27152–27168.
- [15] N. Ibrahim, S. Aboulela, A. Ibrahim, and R. Kashief, "A survey on augmenting knowledge graphs (kgs) with large language models (llms): models, evaluation metrics, benchmarks, and challenges," *Discover Artificial Intelligence*, vol. 4, no. 1, p. 76, 2024.
- [16] Y. Zhou, W. Li, Y. Lu, J. Li, F. Liu, M. Zhang, Y. Wang, D. He, H. Liu, and M. Zhang, "Reflection on knowledge graph for large language models reasoning," in *Findings of ACL*, 2025, pp. 23840–23857.
- [17] Y. Jiang, C. Huang, and L. Huang, "Adaptive graph contrastive learning for recommendation," in *Proc. ACM SIGKDD*, 2023, pp. 4252–4261.
- [18] M. Chen, C. Huang, L. Xia, W. Wei, Y. Xu, and R. Luo, "Heterogeneous graph contrastive learning for recommendation," in *Proc. WSDM*, 2023, pp. 544–552.
- [19] X. Ren, L. Xia, J. Zhao, D. Yin, and C. Huang, "Disentangled contrastive collaborative filtering," in *Proc. ACM SIGIR*, 2023, pp. 1137–1146.
- [20] J. Wu, X. Wang, F. Feng, X. He, L. Chen, J. Lian, and X. Xie, "Self-supervised graph learning for recommendation," in *Proc. ACM SIGIR*, 2021, pp. 726–735.
- [21] J. Yu, H. Yin, X. Xia, T. Chen, L. Cui, and Q. V. H. Nguyen, "Are graph augmentations necessary? simple graph contrastive learning for recommendation," in *Proc. ACM SIGIR*, 2022, pp. 1294–1303.
- [22] Y. Yang, C. Huang, L. Xia, and C. Li, "Knowledge graph contrastive learning for recommendation," in *Proc. ACM SIGIR*, 2022, pp. 1434–1443.
- [23] D. Zou, W. Wei, X.-L. Mao, Z. Wang, M. Qiu, F. Zhu, and X. Cao, "Multi-level cross-view contrastive learning for knowledge-aware recommender system," in *Proc. ACM SIGIR*, 2022, pp. 1358–1368.
- [24] Y. Ma, Y. He, X. Wang, Y. Wei, X. Du, Y. Fu, and T.-S. Chua, "Multicbr: Multi-view contrastive learning for bundle recommendation," *ACM Transactions on Information Systems*, vol. 42, no. 4, pp. 1–23, 2024.
- [25] X. He, K. Deng, X. Wang, Y. Li, Y. Zhang, and M. Wang, "Lightgcn: Simplifying and powering graph convolution network for recommendation," in *Proc. ACM SIGIR*, 2020, pp. 639–648.
- [26] J. Yu, X. Xia, T. Chen, L. Cui, N. Q. V. Hung, and H. Yin, "Xsimgcl: Towards extremely simple graph contrastive learning for recommendation," *IEEE Transactions on Knowledge and Data Engineering*, vol. 36, no. 2, pp. 913–926, 2023.
- [27] Y. Zhang, L. Sang, and Y. Zhang, "Exploring the individuality and collectivity of intents behind interactions for graph collaborative filtering," in *Proc. ACM SIGIR*, 2024, pp. 1253–1262.
- [28] F. Zhang, N. J. Yuan, D. Lian, X. Xie, and W.-Y. Ma, "Collaborative knowledge base embedding for recommender systems," in *Proc. ACM SIGKDD*, 2016, pp. 353–362.
- [29] Y. Zhang, Q. Ai, X. Chen, and P. Wang, "Learning over knowledge-base embeddings for recommendation. arxiv 2018," *arXiv preprint arXiv:1803.06540*, 2018.
- [30] Z. Hu, Z. Li, Z. Jiao, S. Nakagawa, J. Deng, S. Cai, T. Zhou, and F. Ren, "Bridging the user-side knowledge gap in knowledge-aware recommendations with large language models," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 39, no. 11, 2025, pp. 11799–11807.
- [31] X. Ren, W. Wei, L. Xia, L. Su, S. Cheng, J. Wang, D. Yin, and C. Huang, "Representation learning with large language models for recommendation," in *Proc. WWW*, 2024, pp. 3464–3475.