

MLLMRec: A Preference Reasoning Paradigm with Graph Refinement for Multimodal Recommendation

Yuzhuo Dang
dangyuzhuo@nudt.edu.cn
National Key Laboratory of
Information Systems Engineering,
National University of Defense
Technology
Changsha, China

Yuxiao Duan
duanyuxiao19@nudt.edu.cn
Laboratory for Big Data and Decision,
National University of Defense
Technology
Changsha, China

Xin Zhang
zhangxin16@nudt.edu.cn
National Key Laboratory of
Information Systems Engineering,
National University of Defense
Technology
Changsha, China

Wanyu Chen
wanyuchen@nudt.edu.cn
College of Electronic
Countermeasures, National
University of Defense Technology
Hefei, China

Zhiqiang Pan
panzhiqiang@nudt.edu.cn
National Key Laboratory of
Information Systems Engineering,
National University of Defense
Technology
Changsha, China

Fei Cai*
caifei08@nudt.edu.cn
National Key Laboratory of
Information Systems Engineering,
National University of Defense
Technology
Changsha, China

Honghui Chen*
chenhonghui@nudt.edu.cn
National Key Laboratory of
Information Systems Engineering,
National University of Defense
Technology
Changsha, China

Abstract

Multimodal recommendation combines the user historical behaviors with the modal features of items to capture the tangible user preferences, presenting superior performance compared to the conventional ID-based recommender systems. However, existing methods still encounter two key problems in the representation learning of users and items, respectively: (1) the initialization of multimodal user representations is either agnostic to historical behaviors or contaminated by irrelevant modal noise, and (2) the widely used KNN-based item-item graph contains noisy edges with low similarities and lacks audience co-occurrence relationships. To address such issues, we propose MLLMRec, a novel preference reasoning paradigm with graph refinement for multimodal recommendation. Specifically, on the one hand, the item images are first converted into high-quality semantic descriptions using a multimodal large language model (MLLM), thereby bridging the semantic gap between visual and textual modalities. Then, we construct a behavioral description list for each user and feed it into the MLLM to reason about the purified user preference profiles that contain the latent

interaction intents. On the other hand, we develop the threshold-controlled denoising and topology-aware enhancement strategies to refine the suboptimal item-item graph, thereby improving the accuracy of item representation learning. Extensive experiments on three publicly available datasets demonstrate that MLLMRec achieves the state-of-the-art performance with an average improvement of 21.48% over the optimal baselines. The source code is provided at <https://github.com/Yuzhuo-Dang/MLLMRec>.

CCS Concepts

• Information systems → Recommender systems.

Keywords

Multimodal Recommendation, User Preference Reasoning, Item-Item Graph Refinement, Multimodal Large Language Model

ACM Reference Format:

Yuzhuo Dang, Xin Zhang, Zhiqiang Pan, Yuxiao Duan, Wanyu Chen, Fei Cai, and Honghui Chen. 2026. MLLMRec: A Preference Reasoning Paradigm with Graph Refinement for Multimodal Recommendation. In *Proceedings of the 49th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '26)*, July 20–24, 2026, Melbourne, VIC, Australia. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3805712.3809749>

1 Introduction

With the exponential growth of data scale and the diversification of content forms on online platforms, multimodal recommender systems (MMRS) have emerged as a pivotal technology for mitigating

*Corresponding authors.



This work is licensed under a Creative Commons Attribution 4.0 International License. *SIGIR '26, Melbourne, VIC, Australia.*

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2599-9/2026/07

<https://doi.org/10.1145/3805712.3809749>

the information overload [21, 31]. By integrating the various item modalities (e.g., text, image, and audio) with the historical user-item interactions, MMRS facilitate a more nuanced understanding of user preferences, thereby delivering highly personalized services in increasingly complex information environments [15, 28].

The evolution of MMRS has progressed through several distinct paradigms. Pioneering research, exemplified by VBPR [9], focuses on simple feature fusion by augmenting the ID embeddings with the visual features. Subsequently, to capture the high-order collaborative signals, graph-based methods like MMGCN [34] and GRCN [33] introduce graph convolutional networks (GCNs) to propagate the multimodal information across the user-item interaction graph. Recognizing that the item-side correlations are often underutilized, some studies shift their focus toward mining the latent structural patterns among items. For instance, LATTICE [41] constructs the tunable item-item graphs by computing the cosine similarities between items based on the raw and projected modal features. Further, FREEDOM [43] argues that the frozen item-item graphs constructed using only the raw features yield superior performance. Most recently, the burgeoning field of large language models (LLMs) open new frontiers for MMRS [1, 4, 30]. For example, LLMRec [32] leverages the vast world knowledge of LLMs to generate the supplementary attributes for users and items. Moreover, DOGE [20] employs multimodal large language models (MLLMs) to address the issue of granularity inconsistency across modalities, achieving the cross-modal relation mining.

Despite these remarkable achievements, existing MMRS still exhibit the following two challenges in the representation learning at both the user and item sides: (1) **Inaccurate initialization of user representations.** Current approaches typically initialize the user representations through either random initialization or aggregating the modal features of the historically interacted items. However, both paradigms are flawed. Specifically, the former method is inherently behavior-agnostic, neglecting to leverage the rich preference signals in the interactions and thus exacerbating the cold-start issue. The latter approach is often noise-contaminated, as it fails to filter out the modal features encoded by pre-trained models, leading to the indiscriminate aggregation of irrelevant features (e.g., the brown-marked noise in Figure 1(a)). More critically, these methods are incapable of capturing the latent interaction intents hidden in the user decision-making process. Consequently, they are forced to rely on the sophisticated subsequent behavioral modeling to compensate for the informational deficiencies in the initial representations. (2) **Structural noise in the item-item graph topology.** The widely used KNN sparsification strategy [2] for the item-item graph construction is limited by its fixed number of neighbors. This mechanism compels items lacking strongly similar neighbors to associate with low-similarity ones, inevitably introducing false-positive edges. As the evidence reveals in Figure 1(b), approximately 12% of edges in the Baby dataset link items with a similarity score below 0.5. In addition, the KNN strategy depends exclusively on the similarity of modal features while ignoring the latent audience co-occurrence correlations, which are the vital collaborative signals in the field of recommender systems [19]. Consequently, items that are semantically distant but behaviorally correlated remain disconnected (e.g., the basketball and genouillere in Figure 1(c)), resulting in false-negative edges. Both types of structural noise disrupt the

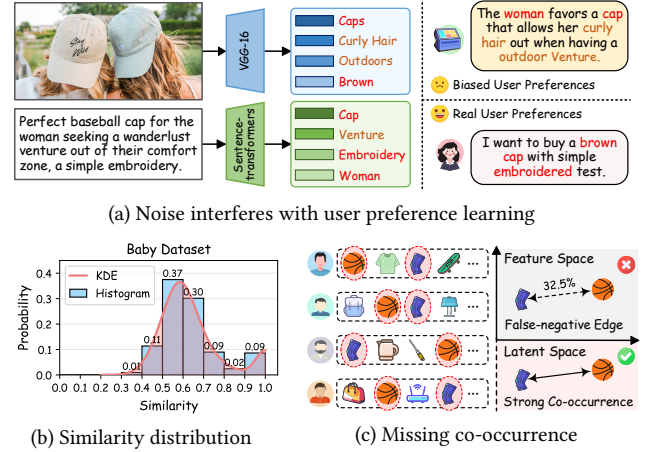


Figure 1: Illustrations of (a) Noise interferes with user preference learning, (b) Distribution of similarity on the Baby dataset, and (c) Missing co-occurrence.

efficacy of the subsequent graph convolution, ultimately leading to semantic drift in the item representation learning.

To address such limitations, we propose a novel preference reasoning paradigm with graph refinement (MLLMRec) for multimodal recommendation. Specifically, to ensure high-quality initialization of user representations, we first utilize the MLLM to convert the item images into semantic descriptions, thereby facilitating the cross-modal alignment. Then, these visual descriptions are integrated with the textual metadata to form the comprehensive multimodal descriptions of items. Distinct from conventional feature extraction, we embed these item-level semantic descriptions into the historical interaction sequences to construct the natural language descriptions of the user behaviors. Building upon this, we leverage the advanced cognitive priors of the MLLM to reason over these descriptions, distilling latent interaction intents while effectively filtering modality-inherent noise to yield the behavior-aware and robust initial user preference profiles. In addition, to rectify the structural noise in the item-item graph, we design two refinement strategies. A threshold-controlled denoising strategy is implemented to prune the spurious edges with low similarity, while a topology-aware enhancement strategy is developed to recover the latent audience co-occurrence correlations. Together, these strategies construct a more reliable item-item graph topology, providing a robust structural foundation for the subsequent graph convolution to capture high-order multimodal features with greater precision. Extensive experiments on three datasets including Baby, Sports, and Clothing validate the effectiveness and robustness of MLLMRec.

Our main contributions can be summarized as follows:

- We propose an innovative multimodal recommendation paradigm that utilizes the MLLM to reason about the user preference profiles, thereby distilling the latent interaction intents and filtering out the modality-inherent noise. Moreover, these initial user representations circumvent reliance on the user-item graph convolution common in standard training pipelines.
- We develop two plug-and-play graph refinement strategies, comprising threshold-controlled denoising and topology-aware enhancement, to eliminate the low-similarity semantic edges and

recover the latent co-occurrence links, respectively, thereby constructing the more reliable item-item graph.

- Extensive experiments conducted on three publicly available datasets demonstrate that MLLMRec outperforms the state-of-the-art baselines, with an average performance improvement of 21.48% over the strongest baseline.

2 Related Works

2.1 Multimodal Recommendation

Unlike conventional recommender systems that rely solely on the user-item interactions, MMRS leverage the rich item-side modalities to facilitate a more comprehensive understanding of user preferences [3, 37]. This field has evolved through several key paradigms. Early efforts, such as VBPR [9], introduces the visual features as auxiliary signals to bolster the collaborative filtering (CF). With the advent of GCNs, researchers shift toward the high-order representation learning. For instance, MMGCN [34] performs message passing on the modality-specific user-item interaction graphs. To further capture the latent item-side correlations, subsequent works like LATTICE [41] and FREEDOM [43] construct the item-item semantic graphs to perform GCNs. However, these methods ignore the different impacts of various modalities on the user decisions. To address this, MGCN [39] incorporates the behavior-aware modal weights, while GUME [16] explicitly models user interests through the extended modality interactions. More recently, self-supervised learning (SSL) has been introduced into MMRS to mitigate the data sparsity and the modal noise. For example, MMSSL [31] employs data augmentation and contrastive learning to mine the inter-modal dependencies. Moreover, MENTOR [36] argues that the excessive modality alignment may lead to the loss of interaction information, proposing a multi-level alignment strategy instead. ModalSync [18] further integrates the supervised and self-supervised paradigms, achieving the synergistic optimization of the behavior modeling and multimodal learning. However, these methods primarily focus on the item feature extraction while resorting to the shallow and incorrect initializations for the user representations. Furthermore, the constructed item-item graph is noisy and incomplete, which impede the learning of robust item representations.

2.2 LLMs-Enhanced Recommendation

With the superior capabilities demonstrated by LLMs in natural language processing tasks, researchers have increasingly explored their potential applications in the recommender systems [17, 29, 40]. Early explorations, such as P5 [6], directly employs a series of prompts to convert various recommendation metadata into a generative sequence-to-sequence framework. Chat-Rec [5] feeds the user behavior sequences into Chat-GPT to generate the personalized recommendations. Subsequently, to align the vast world knowledge of LLMs with the recommendation task, TALLRec [1] uses LoRA [11] for the parameter-efficient fine-tuning of LLaMA, thereby improving its applicability to recommendation. Moreover, ED² [38] proposes a dual dynamic indexing mechanism that enables the synergistic fusion between semantic knowledge and historical interactions. Moreover, to overcome the inherent knowledge limitations of LLMs, RecMind [30] integrates the database querying

and web searching to improve the reasoning accuracy, while introducing the self-inspiring learning for better planning. Furthermore, CIKGR [12] designs a cross-domain contrastive learning framework that aligns the auxiliary information domain with the recommendation task. Recently, several studies pivot toward MLLMs to handle richer content formats. For instance, DOGE [20] employs MLLMs to translate the item images into the descriptive text, effectively unifying disparate the modal granularities. Meanwhile, NEGGEN [13] leverages the generative prowess of MLLMs to synthesize the modality balanced negative samples, thereby providing the robust supervisory signals. However, the existing MLLMs-based methods primarily treat MLLMs as the sophisticated feature translators or data augmentors. They overlook the potential of MLLMs to perform preference reasoning by synthesizing the multimodal contents and the interaction patterns.

3 Task Formulation

Let $\mathcal{U} = \{u\}$ and $\mathcal{I} = \{i\}$ denote the sets of users and items, respectively. Their observed historical interactions are represented by a Boolean matrix $\mathbf{R} = [r_{u,i}] \in \{0, 1\}^{|\mathcal{U}| \times |\mathcal{I}|}$, where $r_{u,i} = 1$ indicates an implicit feedback between user u and item i , and $r_{u,i} = 0$ otherwise. Beyond the collaborative backbone, each item $i \in \mathcal{I}$ is associated with the multimodal attributes $\mathcal{M} = \{v, t\}$, where v and t denote the visual and textual domains, respectively [16, 43]. In addition, we represent the collection of item images as $\mathcal{D}^v = \{i^v\}$ and the corresponding textual descriptions as $\mathcal{D}^t = \{i^t\}$.

Given the interaction matrix $\mathbf{R}$ and the multimodal item metadata $\{\mathcal{D}^v, \mathcal{D}^t\}$, the objective of our work is to learn a predictive function $\mathcal{F} : (\mathbf{R}, \mathcal{D}^v, \mathcal{D}^t) \rightarrow \hat{y}_{ui}$, where $\hat{y}_{ui} \in [0, 1]$ denotes the predicted preference score of user u toward item i . The framework $\mathcal{F}$ aims to capture the synergistic effects between collaborative signals and multimodal semantics. Finally, items are ranked in descending order of their predicted scores $\hat{y}_{ui}$ to generate a personalized top- K recommendation list for each user.

4 Methodology

Figure 2 illustrates the overall framework of our proposed MLLMRec, which synergizes the reasoning stages with the refined item-item graph learning for multimodal recommendation. In the subsequent sections, we elaborate on the design of each component.

4.1 Reasoning Strategy

To bridge the semantic gap across heterogeneous modalities and homogenize the data granularities, we first convert the item images into semantic descriptions. Subsequently, we construct the personalized prompts for each user to guide the MLLM in reasoning about the complex behavioral patterns from the user-item historical interactions, thereby distilling the purified user preference profiles that contain the latent interaction intents.

4.1.1 Semantic Description Generation. Leveraging the sophisticated cross-modal understanding of MLLMs, MLLMRec generates the high-quality semantic descriptions for item images, which is formulated as follows:

$$i^s = \text{MLLM}(\text{Prompt}_1, i^v), \quad (1)$$

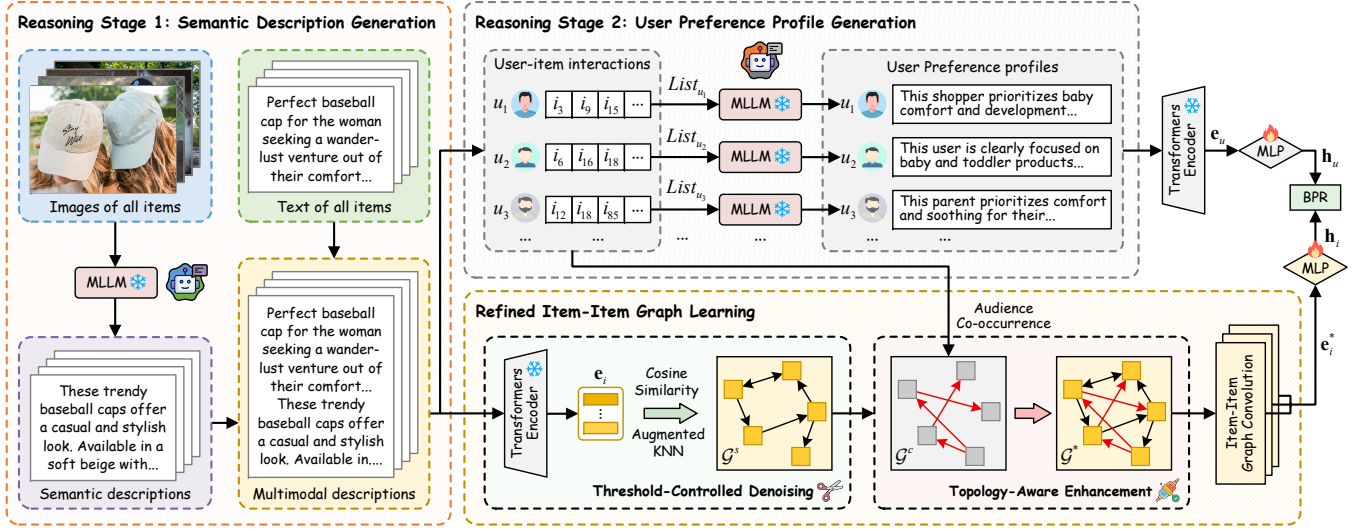


Figure 2: The overall framework of MLLMRec. First, the MLLM transforms the images into the semantic descriptions, which are fused with the text to yield the multimodal descriptions. Next, MLLMRec constructs the behavioral description lists for users, which are fed into the MLLM to reason about the user preference profiles. Meanwhile, MLLMRec optimizes the item-item graph through two designed refinement strategies, subsequently learning high-order item representations on this refined graph.

where $i^s \in \mathcal{D}^s$ denotes the generated semantic description of image $i^v \in \mathcal{D}^v$, and the specific template of Prompt₁ is as follows:

Prompt₁: Semantic Description Generation

{image}: Please convert the given image into an accurate and concise textual description relevant to the {dataset name}, focusing on extracting key attributes that can influence the buying behavior of users, such as color, material, style, functionality, etc. To generate the textual description using a one-paragraph natural language overview in no more than 100 words.

Input: {An Image}

Output: {A paragraph of semantic descriptions for this image}

This stage is pivotal for two reasons. First, conventional visual encoders (e.g., VGG16 [24], ResNet [8]) often capture general features that may be tangential to the specific recommendation contexts. By contrast, MLLMRec designs a task-oriented prompt to guide the MLLM in selectively extracting the decision-relevant visual cues, such as aesthetic style. This yields the semantic descriptions that are aligned with the recommendation task and filter out the noise of visual modality. Second, this generative process projects the visual information into the linguistic space, effectively unifying the data formats of visual and textual modalities. Such homogenization directly addresses the modality imbalance problem, which is a common bottleneck in the multimodal learning where models disproportionately rely on textual data due to the high noise-to-signal ratio and the semantic deficiency of raw visual features [13, 20]. In summary, by transforming the images into the structured semantic description, we ensure that visual signals contribute equitably alongside textual signals to the preference learning.

Subsequently, we integrate i^s with its raw textual metadata i^t to construct a comprehensive multimodal description for each item. The fusion process is formulated as follows:

$$i^m = i^t \oplus i^s, \quad (2)$$

where $i^m \in \mathcal{D}^m$ denotes the multimodal description of item i , and $\oplus$ represents the concatenation operation. By integrating abstract visual descriptions with explicit textual metadata, facilitates the seamless fusion of multimodal information and effectively bridges the cross-modal semantic gap. Consequently, it establishes a purified and enriched multimodal knowledge base that serves as a robust foundation for the subsequent module.

4.1.2 User Preference Profile Generation. Distinguishing from conventional paradigms that rely on the behavior-agnostic random embeddings or the noise-contaminated feature aggregation, MLLMRec capitalizes on the advanced cognitive reasoning of MLLMs to transform the raw interaction data into the purified user preference profiles. These profiles provide a high-fidelity semantic foundation for robust user representation learning. Specifically, we first synthesize a behavioral descriptions list for each user by chronologically aggregating the multimodal descriptions of their historically interacted items. This process is formulated as follows:

$$List_u = [i^m], \quad r_{u,i} = 1, \quad (3)$$

where $List_u$ denotes the behavioral description list of user u , and i^m is the multimodal description of item i that user u has interacted with. Subsequently, we embed $List_u$ into a meticulously designed prompt to guide the MLLM in reasoning about the user's profiles. The generative process is formulated as:

$$u^p = \text{MLLM}(\text{Prompt}_2, List_u), \quad (4)$$

where $u^p \in \mathcal{D}^p$ denotes the generated user preference profiles for user u , and the specific template of Prompt₂ is as follows:

Prompt₂: User Preference Profile Generation

Please reason about the user preferences based on the following list of item descriptions that he or she has interacted with. The list is: $\{behavioral\ description\}$. To generate the user preferences using a one-paragraph natural language in no more than 100 words.

Input: $\{User's\ behavioral\ description\ list\}$

Output: $\{User's\ preference\ profiles\ in\ natural\ language\}$

By utilizing the extensive world knowledge inherent in MLLMs, the resulting user preference profile set $\mathcal{D}^P$ offers a more real user interests compared to the conventional paradigms. It is noteworthy that since $List_u$ is derived from the interaction matrix $\mathbf{R}$, $\mathcal{D}^P$ already captures the latent motivations behind the user-item interactions. This inherent information richness enables the subsequent training models to avoid the need for the explicit structural encoding, such as the user-item graph convolution, thereby reducing computational overhead and streamlining the training process.

4.2 Refined Item-Item Graph Learning

To mitigate the semantic drift in the item representation learning caused by the suboptimal item-item graph topology, we design two graph refinement strategies comprising (1) threshold-controlled denoising and (2) topology-aware enhancement.

4.2.1 Threshold-Controlled Denoising. First, we encode the multimodal descriptions $\mathcal{D}^m$ using a pre-trained text encoder to obtain their latent representations, which is formulated as follows:

$$\mathbf{E}_I = \text{Encoder}(\mathcal{D}^m), \quad (5)$$

where $\mathbf{E}_I = \{\mathbf{e}_i\} \in \mathbb{R}^{|I| \times d_i}$ represents the extracted item multimodal features with dimension d_i , and $\text{Encoder}(\cdot)$ denotes our used text encoder. Then, following the established practice [43], we evaluate the semantic affinities between items by computing the cosine similarity of their multimodal features. This process is formulated as follows:

$$s_{a,b} = \frac{(\mathbf{e}_{i_a})^T \mathbf{e}_{i_b}}{\|\mathbf{e}_{i_a}\| \cdot \|\mathbf{e}_{i_b}\|}, \quad (6)$$

where $s_{a,b}$ denotes the similarity score between item a and item b , $\|\cdot\|$ represents the vector norm computation, and $(\cdot)^T$ is the matrix transpose operation.

While the conventional KNN strategy [2] is widely adopted to sparsify the similarity matrix, they suffer from a rigid fixed-degree constraint. This constraint forces items lacking the genuine semantic neighbors to establish the spurious connections with the low-similarity nodes. As demonstrated by our empirical analysis in Figure 1(b), this phenomenon leads to an obvious presence of the false-positive edges (exceeding 10% in the Baby dataset). To rectify this, we introduce a threshold-controlled denoising strategy that combines the top- K selection with a similarity constraint α . This ensures that only highly similar neighbors are retained for each item. This strategy is formulated as follows:

$$\tilde{s}_{a,b} = \begin{cases} 1, & s_{a,b} \in \text{top-}K_s(s_{a,:}) \text{ and } s_{a,b} \geq \alpha, \\ 0, & \text{otherwise,} \end{cases} \quad (7)$$

where vector $\mathbf{s}_{a,:} \in \mathbb{R}^{|I|}$ represents the similarity scores between item i_a and all items, the hyper-parameters K_s and α control the top- K value and the pruning intensity, respectively. The resulting matrix $\tilde{\mathbf{S}} = [\tilde{s}_{a,b}] \in \{0, 1\}^{|I| \times |I|}$ serves as the adjacency matrix for the denoised item-item semantic affinity graph $\mathcal{G}^s = \{I, \mathcal{E}^s\}$, where I is the set of item nodes and $\mathcal{E}^s = \{\langle i_a, i_b \rangle | \tilde{s}_{a,b} = 1\}$ denotes the edge set. This formulation effectively filters out structural noise while preserving high-confidence semantic correlations.

4.2.2 Topology-Aware Enhancement. We argue that the semantic graph $\mathcal{G}^s$ constructed exclusively from the feature similarities fails to capture the item-item co-occurrence correlations that are essential collaborative signals embedded within the user-item interactions [19]. To bridge this gap, we introduce a topology-aware enhancement strategy to recover these missing behavioral links. Specifically, we quantify the audience co-occurrence between item pairs by computing the Jaccard similarity coefficient over their respective interacted user sets. For any two items i_a and i_b , their co-occurrence score $c_{a,b}$ is formulated as follows:

$$c_{a,b} = J(\mathbf{r}_{:,a}, \mathbf{r}_{:,b}) = \frac{|\mathbf{r}_{:,a} \cap \mathbf{r}_{:,b}|}{|\mathbf{r}_{:,a} \cup \mathbf{r}_{:,b}|}, \quad (8)$$

where $\mathbf{r}_{:,a} \in \mathbb{R}^{|U|}$ and $\mathbf{r}_{:,b} \in \mathbb{R}^{|U|}$ are column vectors of the interaction matrix $\mathbf{R}$, denoting the sets of users who have interacted with item i_a and i_b , respectively. The resulting matrix $\mathbf{C} = [c_{a,b}] \in \mathbb{R}^{|I| \times |I|}$ represents the audience co-occurrence matrix. Given the sparsity of interaction data, most item pairs exhibit limited co-occurrence. Consequently, to ensure structural robustness, we utilize a weighted KNN strategy to retain only the top- K neighbors with the highest co-occurrence scores for each item, which is formulated as follows:

$$\tilde{c}_{a,b} = \begin{cases} c_{a,b}, & c_{a,b} \in \text{top-}K_c(c_{a,:}), \\ 0, & \text{otherwise,} \end{cases} \quad (9)$$

where vector $\mathbf{c}_{a,:} \in \mathbb{R}^{|I|}$ represents the audience co-occurrence scores between item i_a and all items, and hyper-parameter K_c controls the density of co-occurrence connections. This resulting weighted matrix $\tilde{\mathbf{C}} = [\tilde{c}_{a,b}] \in \mathbb{R}^{|I| \times |I|}$ serves as the adjacency matrix for the item-item audience co-occurrence graph $\mathcal{G}^c = \{I, \mathcal{E}^c\}$, where I is the set of item nodes and $\mathcal{E}^c = \{\langle i_a, i_b \rangle | \tilde{c}_{a,b} > 0\}$ denotes the edge set with non-zero co-occurrence scores.

Finally, to establish a holistic representation of item-side relationships, we perform a graph fusion by integrating the semantic affinity graph with the audience co-occurrence graph, which is formulated as follows:

$$\mathbf{S}^* = \tilde{\mathbf{S}} + \tilde{\mathbf{C}}, \quad (10)$$

where $\mathbf{S}^* = [\mathbf{s}_{a,b}^*] \in \mathbb{R}^{|I| \times |I|}$ serves as the adjacency matrix of the topology-enhanced item-item graph $\mathcal{G}^* = \{I, \mathcal{E}^*\}$. The fused edge set $\mathcal{E}^* = \mathcal{E}^s \cup \mathcal{E}^c$ effectively integrates the multimodal semantics with the collaborative signals.

4.2.3 Item-Item Graph Convolution. Building upon the refined graph $\mathcal{G}^*$, we employ LightGCN [10], a simplified and effective variant of GCNs, to perform the neighborhood aggregation. This process enables MLLMRec to capture the high-order correlations among items. Specifically, the graph convolution operation at the

l -th layer is formulated as follows:

$$\mathbf{E}_I^{(l)} = \left((\mathbf{N})^{-\frac{1}{2}} \mathbf{S}^* (\mathbf{N})^{-\frac{1}{2}} \right) \mathbf{E}_I^{(l-1)}, \quad (11)$$

where $\mathbf{E}_I^{(l)} \in \mathbb{R}^{|I| \times d_l}$ represents the item multimodal features at the l -th layer, $\mathbf{E}_I^{(0)} = \mathbf{E}_I$ is the initial embeddings derived from the text encoder (as obtained from Eq. (5)), $\mathbf{N}$ denotes the degree matrix of graph $\mathcal{G}^*$ and $\mathbf{N}_{aa} = \sum_b \mathbf{s}_{a,b}^*$. After that, we aggregate the representations across all hidden layers through a summation function, which is formulated as follows:

$$\mathbf{E}_I^* = \sum_{l=0}^L \mathbf{E}_I^{(l)}, \quad (12)$$

where $\mathbf{E}_I^* = \{\mathbf{e}_i^*\} \in \mathbb{R}^{|I| \times d_l}$ represents the high-order representations of items, and hyper-parameter L controls the depth of the graph convolution.

4.3 Optimization

To align the generated user preference profiles $\mathcal{D}^p$ and the high-order item representations $\mathbf{E}_I^*$ with the recommendation task, we project them into a unified vector space for preference scoring. Specifically, we first leverage the text encoder defined in Eq. (5) to transform the natural language-based user preference profiles $\mathcal{D}^p$ into the vector space, which is formulated as follows:

$$\mathbf{E}_U = \text{Encoder}(\mathcal{D}^p), \quad (13)$$

where $\mathbf{E}_U = \{\mathbf{e}_u\} \in \mathbb{R}^{|\mathcal{U}| \times d_l}$ denotes the extracted user representations. Subsequently, to compress these high-dimensional semantic signals into a compact space suitable for recommendation, we employ multilayer perceptrons (MLPs) for non-linear feature transformation. Taking the user side as an example, the projection is formulated as follows:

$$\mathbf{H}_U = \sigma(\mathbf{E}_U \mathbf{W}_1 + \mathbf{b}_1) \mathbf{W}_2 + \mathbf{b}_2, \quad (14)$$

where $\mathbf{H}_U = \{\mathbf{h}_u\} \in \mathbb{R}^{|\mathcal{U}| \times d}$ represents the final user representations, and $\sigma(\cdot)$ is the LeakyReLU activation function. The weight matrices $\mathbf{W}_1 \in \mathbb{R}^{d_l \times d_1}$, $\mathbf{W}_2 \in \mathbb{R}^{d_1 \times d}$ and bias vectors $\mathbf{b}_1 \in \mathbb{R}^{d_1}$, $\mathbf{b}_2 \in \mathbb{R}^d$ are trainable parameters, where d_1 and d represent the dimensions of the hidden and output layer, respectively. Similarly, the final item representations $\mathbf{H}_I = \{\mathbf{h}_i\} \in \mathbb{R}^{|I| \times d}$ are derived by applying a same projection to the item representations $\mathbf{E}_I^*$.

To facilitate preference learning, we optimize the model parameters via the Bayesian Personalized Ranking (BPR) loss [23]. This objective function encourages the model to assign higher preference scores to the observed interactions than to unobserved ones, which is formulated as follows:

$$\mathcal{L}_{BPR} = - \sum_{(u, i_a, i_b) \in \mathcal{T}} \log(\sigma(\hat{y}_{ui_a} - \hat{y}_{ui_b})), \quad (15)$$

where $\hat{y}_{ui} = \mathbf{h}_u^T \cdot \mathbf{h}_i$ represents the predicted user preference score for user u on item i , and $\mathcal{T} = \{(u, i_a, i_b) \mid r_{u, i_a} = 1, r_{u, i_b} = 0\}$ denotes the set of training triplets, where i_a is an observed item and i_b is a negative sample randomly drawn from the unobserved set. Ultimately, items are ranked in descending order of $\hat{y}_{u, i}$ to generate the top- K recommendation list for each user, where K denotes the length of the recommendation list.

Table 1: Statistics of the experimental datasets.

Datasets	#Users	#Items	#Interactions	Density
Baby	19,445	7,050	160,792	0.117%
Sports	35,598	18,357	296,337	0.045%
Clothing	39,387	23,033	278,677	0.031%

5 Experiments

To validate the effectiveness of MLLMRec, we formulate the following six research questions to guide our experiments:

- **RQ1:** How does the performance of MLLMRec compare to the state-of-the-art multimodal recommendation methods?
- **RQ2:** What are the contributions of the core components and different modalities to the overall performance?
- **RQ3:** To what extent does the recommendation efficacy vary when employing different MLLMs?
- **RQ4:** Can our designed graph refinement strategies be effective in other multimodal recommendation methods?
- **RQ5:** How do the critical hyper-parameter variations influence the performance of MLLMRec?
- **RQ6:** Do the user representations learned by the MLLM effectively capture the valuable features?

5.1 Experiment Setup

5.1.1 Datasets. Following [35, 43], we conduct experiments on three widely used datasets from the Amazon product repository¹: (i) Baby, (ii) Sports and Outdoors, and (iii) Clothing, Shoes, and Jewelry (referred to as Baby, Sports, and Clothing, respectively). These datasets cover diverse product categories, interaction sparsities, and scales of user and items, alongside rich multimodal metadata (i.e., item images and textual descriptions). To ensure the reliability of collaborative signals, we apply the 5-core filtering strategy to filter out the users and items with fewer than five interactions. The statistics of the datasets after filtering are summarized in Table 1. For model training, each dataset is randomly partitioned into training, validation and testing sets with a ratio of 8:1:1.

5.1.2 Evaluation Metrics. We adopt two widely recognized top- K recommendation metrics to quantify model performance: Recall ($R@K$) and Normalized Discounted Cumulative Gain (NDCG, $N@K$). Specifically, Recall measures the fraction of the testing items successfully recommended, while NDCG accounts for ranking quality by assigning higher weights to relevant items positioned earlier in the list. Consistent with prior works, we report the average scores across all users in the testing set under $K \in \{10, 20\}$.

5.1.3 Baselines. To demonstrate the superiority of MLLMRec, we compare it against the following three categories of representative baselines: CF-based recommenders, general multimodal recommenders, and MLLMs-based multimodal recommenders. A brief description of each baseline is provided below:

i) CF-based recommenders:

- **BPR** [23] establishes a foundational Bayesian ranking framework that utilizes the pairwise comparison of items to capture the personalized user preferences from the implicit feedback.

¹<http://jmcauley.ucsd.edu/data/amazon>

- **LightGCN** [10] simplifies the vanilla GCN by omitting the redundant non-linear activations and feature transformations.

ii) general multimodal recommenders:

- **VBPR** [9] integrates the visual features from item images into the BPR model, enabling the capture of visual preferences.
- **LATTICE** [41] constructs modality-aware graphs to mine the latent item-item semantic from multimodal features.
- **SLMRec** [26] employs multimodal data augmentation and contrastive learning to capture the cross-modal patterns.
- **FREEDOM** [43] freezes the item-item graph to improve efficiency and incorporates a degree-sensitive pruning strategy to denoise the user-item interaction graph.
- **LGMRec** [7] combines local topological learning with a global hyper-graph module to capture the high-order dependencies.
- **GUME** [16] utilizes mutual information maximization to align explicit interaction features with extended interest profiles.
- **ModalSync** [18] synchronizes multimodal feature extraction with user behaviors using a staged co-training strategy.
- **COHESION** [35] integrates modality fusion and representation learning via a dual-stage fusion strategy and a composite GCN.

iii) MLLMs-based multimodal recommenders:

- **LLM-CF** [25] introduce LLMs into collaborative filtering through a novel instruction-tuning method and in-context learning.
- **LLMRec** [32] employs LLMs to perform graph augmentation by the user profiling and the item attribute enhancement.
- **DOGE** [20] leverages MLLMs to generate cross-modal representations by aligning the image and text information.
- **NEGGEN** [13] uses MLLMs to construct robust negative samples via tailored prompting, providing better contrastive signals.
- **LIP** [14] utilizes the advanced semantic extraction of MLLMs to encode item attributes into a multimodal interaction graph.

5.1.4 Implementation Details. In the reasoning phase, we employ Gemma3-27b as our used MLLM, accessed via the Ollama framework. For fair comparisons with the existing multimodal baselines, we utilize Sentence-Transformer [22] as the text encoder, and align our setup with the unified MMR framework [42]. Specifically, the training batch size is 2048 and the learning rate is 0.001. All trainable parameters are initialized via the Xavier method and optimized by the Adam optimizer. Following previous settings in [35], we set the output dimension d to 64, the hidden layer dimension d_1 to 256, the semantic neighbor count K_s to 10, and the convolution depth L to 1. For the similarity threshold α , we perform a grid search in $\{0.4, 0.5, 0.6, 0.7\}$, and the co-occurrence neighbor count K_c is ranged in $\{5, 10, 15, 20\}$. In addition, to prevent over-fitting, the training process is terminated if R@20 on the validation set does not improve for 20 consecutive epochs, with a cap training limit of 1,000 epochs. All experiments are implemented using PyTorch and evaluated on a 24GB NVIDIA RTX4090 GPU.

5.2 Overall Performance (RQ1)

Table 2 presents the experimental results of the baselines and MLLMRec on three datasets, with the following key findings:

(1) **Multimodal signals enhance the preference modeling.** Integrating the modal features substantially improve the performance of CF-based recommenders. A direct comparison between

BPR and its modality extension VBPR reveals that the latter achieves superior results across all datasets. This confirms that relying solely on the collaborative signals is insufficient to capture the fine-grained user preferences due to the interaction sparsity. Multimodal information, such as visual cues, provides a critical semantic supplement, effectively enriching the representation space.

(2) **MLLMs promote the multimodal feature learning.** Compared to the general multimodal recommenders, MLLMs-based methods generally exhibit stronger competitiveness. This superiority is attributed to the exceptional cognitive reasoning and cross-modal alignment abilities of MLLMs, which facilitate a paradigm shift from shallow the feature extraction to the deep semantic understanding. By transforming the raw data into high-level descriptors, these models better capture the core features, leading to more precise matching in complex recommendation tasks.

(3) **MLLMRec obtains significant performance gains.** Specifically, MLLMRec achieves relative improvements of 27.57%, 16.39%, and 17.50% in R@20, and 30.15%, 17.94%, and 18.13% in N@20 over the optimal baselines on the Baby, Sports, and Clothing datasets, respectively. These results confirm that utilizing MLLMs to reason about the user preference profiles, combined with the item-item graph refinement strategies, effectively enhances the recommendation performance. Furthermore, the statistical significance of these gains is validated via t-tests under five independent runs.

(4) **MLLMRec avoid the need for the user-item graph convolution.** While most baselines rely on GCNs over the user-item graph to fill in the information shortage in the randomly initialized user representations, the comparison between MLLMRec and its variant MLLMRec w/ GCN_{UI} reveals that adding a GCN module actually degrades performance. This suggests that the initial user representations reasoned by MLLMRec already encapsulate the latent interaction intents. Consequently, introducing an extra GCN_{UI} becomes counterproductive, as it induces feature homogenization and potentially introduces structural noise.

5.3 Ablation Study (RQ2)

To evaluate the performance contributions of each core component and different modalities, we design the following six model variants:

(1) *w/o RS*, which replaces the user preference profile generation with the randomly initialized user representations, (2) *w/o GD*, which excludes the threshold-controlled denoising strategy, (3) *w/o TE*, which abandons the topology-aware enhancement strategy, and (4) *w/o GCN_{II}*, which removes the item-item graph convolution, (5) *w/o V* and (6) *w/o T*, which omit the visual and textual modalities, respectively. Table 3 lists the experimental results, and it reveals the following three observations:

(1) **The absence of any core component leads to performance degradation.** Specifically, RS provides the high-fidelity user preference profiles that contain latent interaction intents. GD and TE strategies effectively refine the traditional item-item graph, and GCN_{II} module further leverages this refined graph to extract the high-order representations of items. Therefore, these components are essential for the optimal performance of MLLMRec, and excluding any of them would degrade the performance.

(2) **Removing RS causes the largest decrease in the model performance.** This observation validates the efficacy of employing

Table 2: Overall performance of the baselines and MLLMRec on three datasets, with the best and second-best results in each column highlighted by bold and underline, respectively. *Improv.* represents the percentage improvement of MLLMRec over the optimal baselines. * indicates that the improvements are statistically significant based on paired t-test (p -value < 0.05).

Models	Baby				Sports				Clothing			
	R@10	R@20	N@10	N@20	R@10	R@20	N@10	N@20	R@10	R@20	N@10	N@20
BPR (UAI'09)	0.0357	0.0575	0.0192	0.0249	0.0432	0.0653	0.0241	0.0298	0.0187	0.0279	0.0103	0.0126
LightGCN (SIGIR'20)	0.0479	0.0754	0.0257	0.0328	0.0569	0.0864	0.0311	0.0387	0.0340	0.0526	0.0188	0.0236
VBPR (AAAI'16)	0.0423	0.0663	0.0223	0.0284	0.0558	0.0856	0.0307	0.0384	0.0281	0.0415	0.0158	0.0192
LATTICE (MM'21)	0.0547	0.0850	0.0292	0.0370	0.0620	0.0953	0.0335	0.0421	0.0492	0.0733	0.0268	0.0330
SLMRec (TMM'23)	0.0540	0.0810	0.0285	0.0357	0.0676	0.1017	0.0374	0.0462	0.0452	0.0675	0.0247	0.0303
FREEDOM (MM'23)	0.0627	0.0992	0.0330	0.0424	0.0717	0.1089	0.0385	0.0481	0.0629	0.0941	0.0341	0.0420
LGMRc (AAAI'24)	0.0644	0.1002	0.0349	0.0440	0.0720	0.1068	0.0390	0.0480	0.0555	0.0828	0.0302	0.0371
GUME (CIKM'24)	0.0673	0.1042	0.0365	0.0460	0.0778	0.1165	0.0427	0.0527	0.0703	0.1024	0.0384	0.0466
ModalSync (WWW'25)	0.0693	0.1045	0.0373	0.0462	0.0821	0.1189	0.0457	0.0551	0.0738	0.1058	0.0406	0.0488
COHESION (SIGIR'25)	0.0680	0.1052	0.0354	0.0454	0.0752	0.1137	0.0409	0.0503	0.0665	0.0983	0.0358	0.0438
LLM-CF (CIKM'24)	0.0593	0.0874	0.0312	0.0409	0.0698	0.1057	0.0376	0.0465	0.0596	0.0886	0.0329	0.0402
LLMRec (WSDM'24)	0.0669	0.1037	0.0363	0.0456	0.0826	0.1189	0.0453	0.0554	0.0695	0.1018	0.0377	0.0452
DOGE (AAAI'25)	0.0718	0.1096	0.0391	0.0482	0.0792	0.1185	0.0435	0.0536	0.0712	0.1036	0.0389	0.0468
NEGGEN (MM'25)	0.0701	0.1065	0.0342	0.0438	0.0763	0.1114	0.0411	0.0506	0.0654	0.0961	0.0350	0.0441
LIP (INFFUS'26)	<u>0.0978</u>	<u>0.1422</u>	<u>0.0535</u>	<u>0.0650</u>	<u>0.0849</u>	<u>0.1287</u>	<u>0.0458</u>	<u>0.0574</u>	<u>0.0751</u>	<u>0.1126</u>	<u>0.0399</u>	<u>0.0491</u>
MLLMRec w/ GCN _{UI}	0.1091	0.1582	0.0605	0.0730	0.0870	0.1323	0.0464	0.0580	0.0737	0.1107	0.0393	0.0486
MLLMRec (Ours)	0.1240*	0.1814*	0.0700*	0.0846*	0.1010*	0.1498*	0.0553*	0.0677*	0.0871*	0.1323*	0.0466*	0.0580*
<i>Improv.</i>	26.79%	27.57%	30.84%	30.15%	18.96%	16.39%	20.74%	17.94%	15.98%	17.50%	16.79%	18.13%

Table 3: Performance of MLLMRec and its variants in terms of R@20 and N@20 on three datasets.

Model Variant	Baby		Sports		Clothing	
	R@20	N@20	R@20	N@20	R@20	N@20
MLLMRec	0.1814	0.0846	0.1498	0.0677	0.1323	0.0580
w/o RS	0.0912	0.0416	0.0979	0.0437	0.0746	0.0339
w/o GD	0.1781	0.0838	0.1441	0.0655	0.1302	0.0577
w/o TE	0.1699	0.0804	0.1417	0.0638	0.1246	0.0552
w/o GCN _{II}	0.1672	0.0786	0.1394	0.0631	0.1187	0.0528
w/o V	0.1538	0.0721	0.1233	0.0562	0.1056	0.0457
w/o T	0.1462	0.0686	0.1198	0.0545	0.1012	0.0442

MLLMs to reason about the behavior-aware and robust user preference profiles, which serve as superior initial states compared to the traditional initialization methods. Furthermore, this success highlights a paradigm shift in the field, extending the boundary of multimodal recommendation from the conventional feature engineering to the optimization of the user preference profile generation.

(3) The omission of any modality leads to performance decline. This indicates that MLLMRec effectively leverages the complementary attributes of disparate modalities to capture the multifaceted user interests. Furthermore, by unifying the item images into the high-quality semantic descriptions, MLLMRec mitigates the common issues of modality imbalance and noise interference in the multimodal learning, thereby facilitating a more effective utilization of the multimodal information.

5.4 Compatibility Study (RQ3)

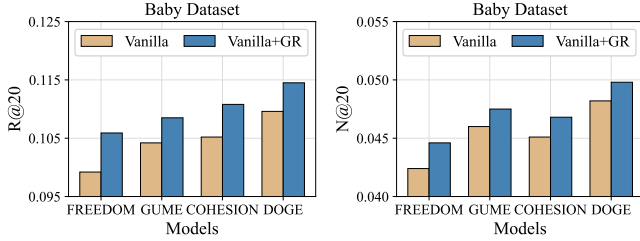
To further investigate the robustness of MLLMRec with respect to the size and type of MLLMs, we select six representative models, including the Gemma3 series (4b, 12b, 27b), Llava-7b, Ministral3-8b, and Qwen2.5vl-7b. The experimental results are shown in Table 4. We can observe that MLLMRec maintains robust performance when using different MLLMs, which demonstrating its efficacy is generalized rather than dependent on the deviation of a specific model. In addition, the incremental performance in the Gemma3 series reveals a clear scaling effect, where stronger logical reasoning capabilities can more accurately capture the visual information and the user preference profiles. Notably, Qwen2.5vl-7b achieves the state-of-the-art results across all datasets, owing to its superior vision-language alignment. This not only confirms the decisive role of the high-quality semantic priors in improving the recommendation accuracy, but also positions MLLMRec as a forward-looking framework capable of capitalizing on the rapid evolution of MLLMs to achieve the iterative performance breakthroughs.

5.5 Plug-and-Play Evaluation (RQ4)

Our proposed item-item graph refinement strategies are designed as plug-and-play components compatible with various multimodal recommendation models that leverage item-item graph learning. To evaluate its transferability, we select four representative models, FREEDOM, GUME, COHESION, and DOGE. For each baseline, we construct an enhanced variant (collectively denoted as Vanilla+GR) by refining their original item-item graphs. As shown in Figure 3, the integration of the GR component yields obvious performance

Table 4: Performance of MLLMRec when using different multimodal large language models. The best results are highlighted in bold. The model adopted in our work and its corresponding results are marked with gray shading.

MLLMs	Baby				Sports				Clothing			
	R@10	R@20	N@10	N@20	R@10	R@20	N@10	N@20	R@10	R@20	N@10	N@20
Gemma3-4b	0.1158	0.1624	0.0642	0.0776	0.0992	0.1453	0.0542	0.0665	0.0824	0.1236	0.0450	0.0554
Gemma3-12b	0.1222	0.1736	0.0671	0.0802	0.1008	0.1487	0.0549	0.0673	0.0826	0.1247	0.0451	0.0557
Gemma3-27b	0.1240	0.1814	0.0700	0.0846	0.1010	0.1498	0.0553	0.0677	0.0871	0.1323	0.0466	0.0580
Llava-7b	0.0864	0.1269	0.0471	0.0574	0.0708	0.1055	0.0387	0.0476	0.0587	0.0882	0.0317	0.0392
Ministral3-8b	0.0943	0.1381	0.0534	0.0645	0.0918	0.1375	0.0508	0.0624	0.0621	0.0951	0.0335	0.0418
Qwen2.5vl-7b	0.1535	0.2092	0.0894	0.1036	0.1100	0.1602	0.0614	0.0742	0.0882	0.1297	0.0485	0.0590

**Figure 3: Effect of plugging the graph refinement strategies into other models on the Baby dataset.**

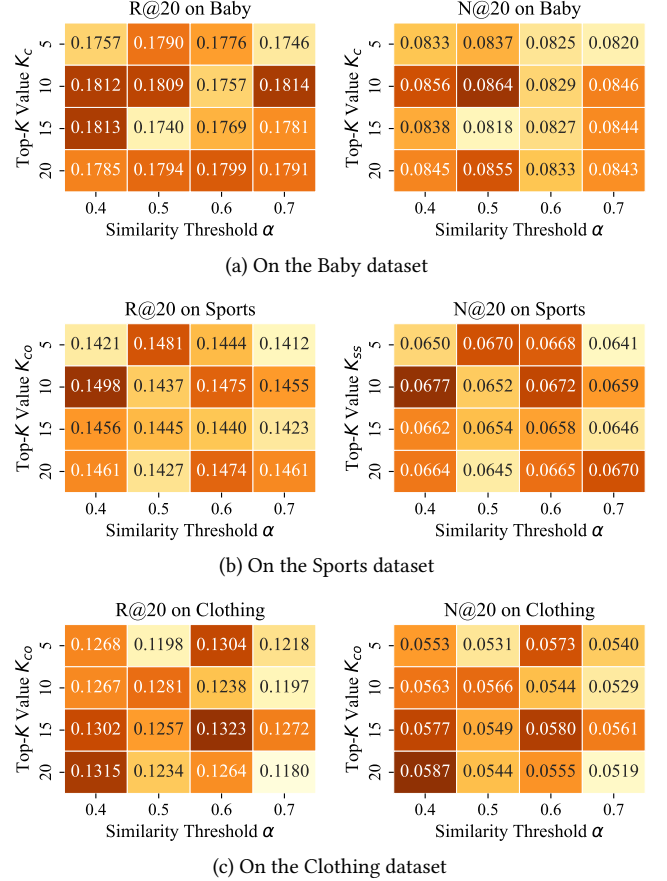
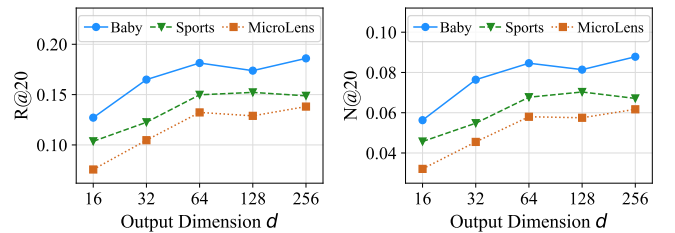
gains across all evaluated models on the Baby dataset. Specifically, the GR-enhanced variants achieve average improvements of 5.17% in R@20 and 3.88% in N@20 compared to their vanilla models. These results provide strong evidence for the robustness and universality of our refinement strategies. By effectively prune the low-similarity links and recover the missing behavioral correlations, the GR strategies effectively improve the quality of the item-item graphs, offering a scalable solution for more reliable message passing.

5.6 Sensitivity Analysis (RQ5)

In this section, we analyze the impact of hyper-parameter variations on MLLMRec, specifically focusing on the similarity threshold α , the co-occurrence top-K value K_c , and the output dimension d .

5.6.1 Impact of α and K_c . We vary the similarity threshold α in $\{0.4, 0.5, 0.6, 0.7\}$ and adjust the top-K value K_c from 5 to 20 in steps of 5. The heatmaps in Figure 4 illustrate the performance fluctuations under different value combinations across three datasets. We observe that the optimal value combination varies across different datasets, which suggests the necessity for careful parameter tuning in practical applications. Specifically, for the clothing dataset, which possesses a larger item pool, a higher K_c is required to capture the diverse behavioral correlations. Furthermore, MLLMRec exhibits relatively lower sensitivity to the variations in K_c compared to α . Notably, while a higher α effectively prunes the structural noise, an excessively high threshold leads to discard the legitimate edges, thereby hindering the effective message passing.

5.6.2 Impact of d . In the prior experiments, we fix the output dimension d at 64 to ensure consistency with the baselines. To further explore the influence of dimension size on the performance of MLLMRec, we tune d in $\{16, 32, 64, 128, 256\}$. As shown in Figure 5, the results indicate that the performance initially exhibits an

**Figure 4: Performance of MLLMRec under different values of α and K_c on three datasets.****Figure 5: Performance under different values of d .**

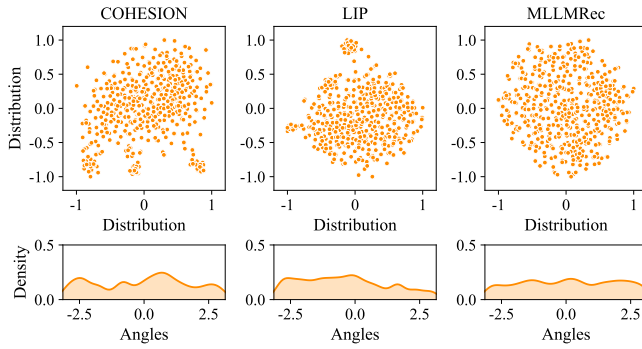


Figure 6: The distribution of the user representations learned by DOGE, ModalSync, and MLLMRec on the Baby dataset. The upper figures depict the 2D feature distributions, and the lower figures show the density of points along the angle.

upward trend as d increasing. This is attributed to the fact that the higher-dimensional spaces offer greater representational capacity to retain the semantic features distilled from MLLMs. However, the performance improvement begin to fluctuate once d exceeds 64. This suggests that beyond a certain threshold, increasing the dimension may introduce the feature redundancy or noise.

5.7 Visualization Analysis (RQ6)

To further validate the superiority of the MLLM-driven user preference profile generation paradigm, we conduct a comparative analysis of the user representation distributions learned by COHESION, LIP, and MLLMRec. Specifically, we randomly sample 500 users from the Baby dataset and project their learned representations into a 2D space using the t-SNE method [27]. As illustrated in Figure 6, the user representations learned by COHESION and LIP exhibit the localized clustering. In contrast, MLLMRec achieves a more decentralized and uniform distribution. According to the recent theoretical insights in the multimodal representation learning [3, 39], such uniformity indicates that the model is effectively preserving a broader spectrum of the semantic features. This finding indicates why the user representations learned by MLLMRec are superior, thereby demonstrating the efficacy of the MLLM-driven paradigm in reasoning about the user preference profiles.

6 Conclusion and Future Work

In this paper, we propose MLLMRec, a novel framework that redefines multimodal recommendation by using the cognitive reasoning of MLLMs and graph structural refinement. Specifically, we introduce a pioneering paradigm that leverages MLLMs to reason about high-fidelity user preference profiles, effectively capturing the latent interaction intents behind the user behaviors. Our findings demonstrate that such superior semantic initialization can avoid the need for the user-item graph convolution. Furthermore, to address the structural noise in the item-item graph, we design the threshold-controlled denoising and topology-aware enhancement strategies to refine the suboptimal item-item graph, which can also be plug-and-play for other models employing the item-item graph learning. Extensive experiments across three publicly datasets demonstrate the superiority and robustness of MLLMRec.

In the future, we aim to investigate temporally aware reasoning strategies to capture the evolving dynamics of user interests, while exploring the use of implicit reasoning tokens to bypass explicit textual steps, thereby improving the computational efficiency.

Acknowledgments

The authors would like to thank the support by the COSTA: complex system optimization team of the College of System Engineering at NUDT. This research was partially supported by the National Natural Science Foundation of China under No.: 62302511 and No.: 62372458. All content represents the opinion of the authors, which is not necessarily shared or endorsed by their respective employers and/or sponsors.

References

- [1] Keqin Bao, Jizhi Zhang, Yang Zhang, Wenjie Wang, Fuli Feng, and Xiangnan He. 2023. TALLRec: An Effective and Efficient Tuning Framework to Align Large Language Model with Recommendation. In *Proceedings of the 17th ACM Conference on Recommender Systems*. 1007–1014.
- [2] Jie Chen, Haw-ren Fang, and Yousef Saad. 2009. Fast Approximate k NN Graph Construction for High Dimensional Data via Recursive Lanczos Bisection. *Journal of Machine Learning Research* 10 (2009), 1989–2012.
- [3] Yuzhuo Dang, Zhiqiang Pan, Xin Zhang, Wanyu Chen, Fei Cai, and Honghui Chen. 2025. Discrepancy Learning Guided Hierarchical Fusion Network for Multi-modal Recommendation. *Knowledge-Based Systems* 317 (2025), 113496.
- [4] Yuxiao Duan, Xiang Zhao, and Hao Guo. 2025. Sentiment-aware cross-modal semantic interaction model for harmful meme detection. *Decision Support Systems* 196 (2025), 114509.
- [5] Yunfan Gao, Tao Sheng, Youlin Xiang, Yun Xiong, Haofen Wang, and Jiawei Zhang. 2023. Chat-REC: Towards Interactive and Explainable LLMs-Augmented Recommender System. *arXiv preprint arXiv:2303.14524* (2023).
- [6] Shijie Geng, Shuchang Liu, Zuohui Fu, Yingqiang Ge, and Yongfeng Zhang. 2022. Recommendation as Language Processing (RLP): A Unified Pretrain, Personalized Prompt & Predict Paradigm (P5). In *Proceedings of the 16th ACM Conference on Recommender Systems*. 299–315.
- [7] Zhiqiang Guo, Jianjun Li, Guohui Li, Chaoyang Wang, Si Shi, and Bin Ruan. 2024. LGMRec: Local and Global Graph Learning for Multimodal Recommendation. In *Proceedings of the 38th AAAI Conference on Artificial Intelligence*. 8454–8462.
- [8] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the 29th IEEE conference on computer vision and pattern recognition*. 770–778.
- [9] Ruining He and Julian J. McAuley. 2016. VBPR: Visual Bayesian Personalized Ranking from Implicit Feedback. In *Proceedings of the 30th AAAI Conference on Artificial Intelligence*. 144–150.
- [10] Xiangnan He, Kuan Deng, Xiang Wang, Yan Li, Yong-Dong Zhang, and Meng Wang. 2020. LightGCN: Simplifying and Powering Graph Convolution Network for Recommendation. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*. 639–648.
- [11] Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2022. LoRA: Low-Rank Adaptation of Large Language Models. In *Proceedings of the 10th International Conference on Learning Representations*.
- [12] Zheng Hu, Zhe Li, Ziyun Jiao, Satoshi Nakagawa, Jiawen Deng, Shimin Cai, Tao Zhou, and Fuji Ren. 2025. Bridging the User-side Knowledge Gap in Knowledge-aware Recommendations with Large Language Models. In *Proceedings of the 39th AAAI Conference on Artificial Intelligence*. 11799–11807.
- [13] Yanbiao Ji, Dan Luo, Chang Liu, Shaokai Wu, Jing Tong, Qichen He, Deyi Ji, Hongtao Lu, and Yue Ding. 2025. Generating Negative Samples for Multi-Modal Recommendation. In *Proceedings of the 33rd ACM International Conference on Multimedia*. 6007–6016.
- [14] Meng Jian, Ruoxi Li, Tuo Wang, Ge Shi, and Lifang Wu. 2026. Large multimodal model-based intent prompting learning for multimedia recommendation. *Information Fusion* 127 (2026), 103881.
- [15] Zhenyang Li, Fan Liu, Yinwei Wei, Zhiyong Cheng, Liqiang Nie, and Mohan S. Kankanhalli. 2024. Attribute-driven Disentangled Representation Learning for Multimodal Recommendation. In *Proceedings of the 32nd ACM International Conference on Multimedia*. 9660–9669.
- [16] Guojiao Lin, Zhen Meng, Dongjie Wang, Qingqing Long, Yuanchun Zhou, and Meng Xiao. 2024. GUME: Graphs and User Modalities Enhancement for Long-Tail Multimodal Recommendation. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management*. 1400–1409.

- [17] Jianghao Lin, Xinyi Dai, Yunjia Xi, Weiwen Liu, Bo Chen, Hao Zhang, Yong Liu, Chuhan Wu, Xiangyang Li, Chenxu Zhu, Huifeng Guo, Yong Yu, Ruiming Tang, and Weinan Zhang. 2025. How Can Recommender Systems Benefit from Large Language Models: A Survey. *ACM Transactions on Information Systems* 43, 2 (2025), 28:1–28:47.
- [18] Shiqin Liu, Chaozhao Li, Minjun Zhao, Litian Zhang, and Jiajun Bu. 2025. Modal-Sync: Synchronizing User Behavior with Multimodal Features for Multimodal Pre-training Recommendation. In *Proceedings of the ACM Web Conference 2025*. 2278–2287.
- [19] Yaokun Liu, Xiaowang Zhang, Minghui Zou, and Zhiyong Feng. 2023. Co-occurrence embedding enhancement for long-tail problem in multi-interest recommendation. In *Proceedings of the 17th ACM Conference on Recommender Systems*. 820–825.
- [20] Fanshen Meng, Zhenhua Meng, Ru Jin, Rongheng Lin, and Budan Wu. 2025. DOGE: LLMs-Enhanced Hyper-Knowledge Graph Recommender for Multimodal Recommendation. In *Proceedings of the 39th AAAI Conference on Artificial Intelligence*. 12399–12407.
- [21] Zongshen Mu, Yueting Zhuang, Jie Tan, Jun Xiao, and Siliang Tang. 2022. Learning Hybrid Behavior Patterns for Multimedia Recommendation. In *Proceedings of the 30th ACM International Conference on Multimedia*. 376–384.
- [22] Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing*. 3980–3990.
- [23] Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. 2009. BPR: Bayesian Personalized Ranking from Implicit Feedback. In *Proceedings of the 25th Conference on Uncertainty in Artificial Intelligence*. 452–461.
- [24] Karen Simonyan and Andrew Zisserman. 2014. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556* (2014).
- [25] Zhongxiang Sun, Zihua Si, Xiaoxue Zang, Kai Zheng, Yang Song, Xiao Zhang, and Jun Xu. 2024. Large language models enhanced collaborative filtering. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management*. 2178–2188.
- [26] Zhulin Tao, Xiaohao Liu, Yewei Xia, Xiang Wang, Lifang Yang, Xianglin Huang, and Tat-Seng Chua. 2023. Self-Supervised Learning for Multimedia Recommendation. *IEEE Transactions on Multimedia* 25 (2023), 5107–5116.
- [27] Laurens Van der Maaten and Geoffrey Hinton. 2008. Visualizing data using t-SNE. *Journal of Machine Learning Research* 9, 11 (2008).
- [28] Qifan Wang, Yinwei Wei, Jianhua Yin, Jianlong Wu, Xuemeng Song, and Liqiang Nie. 2023. DualGNN: Dual Graph Neural Network for Multimedia Recommendation. *IEEE Transactions on Multimedia* 25 (2023), 1074–1084.
- [29] Yan Wang, Zhixuan Chu, Xin Ouyang, Simeng Wang, Hongyan Hao, Yue Shen, Jinjie Gu, Siqiao Xue, James Zhang, Qing Cui, Longfei Li, Jun Zhou, and Sheng Li. 2024. LLMRG: Improving Recommendations through Large Language Model Reasoning Graphs. In *Proceedings of the 38th AAAI Conference on Artificial Intelligence*. 19189–19196.
- [30] Yancheng Wang, Ziyang Jiang, Zheng Chen, Fan Yang, Yingxue Zhou, Eunah Cho, Xing Fan, Yanbin Lu, Xiaojiang Huang, and Yingzhen Yang. 2024. RecMind: Large Language Model Powered Agent For Recommendation. In *Findings of the Association for Computational Linguistics*. 4351–4364.
- [31] Wei Wei, Chao Huang, Lianghao Xia, and Chuxu Zhang. 2023. Multi-Modal Self-Supervised Learning for Recommendation. In *Proceedings of the ACM Web Conference 2023*. 790–800.
- [32] Wei Wei, Xubin Ren, Jiabin Tang, Qinyong Wang, Lixin Su, Suqi Cheng, Junfeng Wang, Dawei Yin, and Chao Huang. 2024. LLMRec: Large Language Models with Graph Augmentation for Recommendation. In *Proceedings of the 17th ACM International Conference on Web Search and Data Mining*. 806–815.
- [33] Yinwei Wei, Xiang Wang, Liqiang Nie, Xiangnan He, and Tat-Seng Chua. 2020. Graph-Refined Convolutional Network for Multimedia Recommendation with Implicit Feedback. In *Proceedings of the 28th ACM International Conference on Multimedia*. 3541–3549.
- [34] Yinwei Wei, Xiang Wang, Liqiang Nie, Xiangnan He, Richang Hong, and Tat-Seng Chua. 2019. MMGCN: Multi-modal Graph Convolution Network for Personalized Recommendation of Micro-video. In *Proceedings of the 27th ACM International Conference on Multimedia*. 1437–1445.
- [35] Jinfeng Xu, Zheyu Chen, Wei Wang, Xiping Hu, Sang-Wook Kim, and Edith CH Ngai. 2025. COHESION: Composite graph convolutional network with dual-stage fusion for multimodal recommendation. In *Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 1830–1839.
- [36] Jinfeng Xu, Zheyu Chen, Shuo Yang, Jinze Li, Hewei Wang, and Edith C. H. Ngai. 2025. MENTOR: Multi-level Self-supervised Learning for Multimodal Recommendation. In *Proceedings of the 39th AAAI Conference on Artificial Intelligence*. 12908–12917.
- [37] Zixuan Yi, Xi Wang, Iadh Ounis, and Craig MacDonald. 2022. Multi-modal Graph Contrastive Learning for Micro-video Recommendation. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 1807–1811.
- [38] Jun Yin, Zhengxin Zeng, Mingzheng Li, Hao Yan, Chaozhao Li, Weihao Han, Jianjin Zhang, Ruochen Liu, Hao Sun, Weiwei Deng, Feng Sun, Qi Zhang, Shirui Pan, and Senzhang Wang. 2025. Unleash LLMs Potential for Sequential Recommendation by Coordinating Dual Dynamic Index Mechanism. In *Proceedings of the ACM on Web Conference 2025*. 216–227.
- [39] Penghang Yu, Zhiyi Tan, Guanming Lu, and Bing-Kun Bao. 2023. Multi-View Graph Convolutional Network for Multimedia Recommendation. In *Proceedings of the 31st ACM International Conference on Multimedia*. 6576–6585.
- [40] Zhenrui Yue, Sara Rabhi, Gabriel de Souza Pereira Moreira, Dong Wang, and Even Oldridge. 2023. LlamaRec: Two-Stage Recommendation using Large Language Models for Ranking. *arXiv preprint arXiv:2311.02089* (2023).
- [41] Jinghao Zhang, Yanqiao Zhu, Qiang Liu, Shu Wu, Shuhui Wang, and Liang Wang. 2021. Mining Latent Structures for Multimedia Recommendation. In *Proceedings of the 29th ACM International Conference on Multimedia*. 3872–3880.
- [42] Xin Zhou. 2023. MMRec: Simplifying Multimodal Recommendation. In *ACM Multimedia Asia Workshops*. 6:1–6:2.
- [43] Xin Zhou and Zhiqi Shen. 2023. A Tale of Two Graphs: Freezing and Denoising Graph Structures for Multimodal Recommendation. In *Proceedings of the 31st ACM International Conference on Multimedia*. 935–943.