

Band Together: Untargeted Adversarial Training with Multimodal Coordination against Evasion-based Promotion Attacks

Guanmeng Xian¹, Ning Yang^{1*}, Philip S. Yu²

¹Sichuan University, Chengdu, China

²University of Illinois at Chicago, USA

xianguanmeng@stu.scu.edu.cn, yangning@scu.edu.cn, psyu@uic.edu

Abstract

Multimodal recommender systems exploit visual and textual signals to alleviate data sparsity, but this also makes them more vulnerable to evasion-based promotion attacks. Existing defenses are largely limited to single-modal settings and mainly focus on poisoning-based threats, leaving evasion-based threats underexplored. In this work, we first identify a cross-modal gradient mismatch under the multi-user promotion setting, where visual and textual perturbations are optimized in inconsistent directions due to the dominance of distinct user groups. This phenomenon dilutes the attack effectiveness and leads robust training to underestimate worst-case risks. To address this issue, we propose Untargeted Adversarial Training with Multimodal Coordination (UAT-MC). UAT-MC tackles the challenge of unknown targeted items in evasion-based attacks (as opposed to poisoning-based attacks) by treating all items as potential targets, and introduces a gradient alignment mechanism to explicitly correct this mismatch. This design ensures synchronized perturbations across modalities, thereby maximizing adversarial strength for robust training. Extensive experiments demonstrate that UAT-MC significantly improves robustness against promotion attacks while maintaining acceptable recommendation performance under the defense-accuracy trade-off. Code is available at <https://github.com/gmXian/UAT-MC>.

1 Introduction

Multimodal Recommender Systems (MRSs) typically leverage visual and textual information from user-interacted items to capture users’ fine-grained preferences, effectively alleviating the challenge posed by sparse interaction data [He and McAuley, 2016b; Wei *et al.*, 2019; Zhang *et al.*, 2021; Zhou *et al.*, 2023c; Liu *et al.*, 2024]. While auxiliary modal information enhances the personalization of recommendations, we find that MRSs exhibit heightened vulnerability to

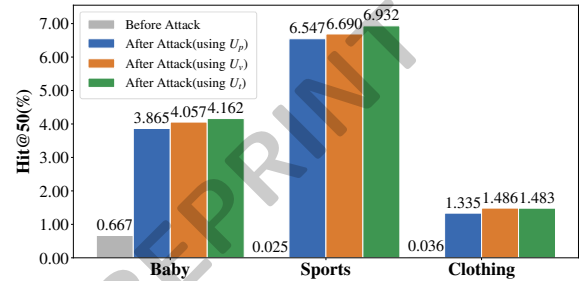


Figure 1: VBPR’s vulnerability to vanilla FGSM-based promotion attacks. For each dataset, we report the Hit@50 before and after attack under different user subsets, including U_p , U_v , and U_t .

evasion-based promotion attacks, in which attackers imperceptibly perturb the modalities of targeted items to boost their rankings, as demonstrated by our motivation experiment in Figure 1.

Specifically, we conduct evasion-based promotion attacks (FGSM [Goodfellow *et al.*, 2014] and PGD [Madry *et al.*, 2017]) on both visual and textual modalities of unpopular items (defined as those with only 5 interactions) over various datasets (Amazon Baby, Sports and Clothing [He and McAuley, 2016a]) and architectures (VBPR [He and McAuley, 2016b] and MMGCN [Wei *et al.*, 2019])). Figure 1 presents the results on VBPR under FGSM attack as a representative instance. The black numbers represent the average hit rate (Hit@50) of targeted items appearing in user recommendation lists. As one can find from Figure 1, the post-attack Hit@50 always shows a significant increase, demonstrating the vulnerability of MRS to evasion-based promotional attacks.

Existing works [Wu *et al.*, 2021; Zhang *et al.*, 2024; Mu *et al.*, 2025] have attempted to mitigate the vulnerability to promotion attacks; however, they suffer from two critical limitations when applied to MRSs. **First**, these approaches primarily focus on single-modal recommender systems, failing to address the complex vulnerability patterns that arise from interactions between visual and textual features. This oversight is especially problematic as attackers can exploit cross-modal correlations to amplify their attack impact. For example, the Re-writing Defense [Zhang *et al.*, 2024] uses GPT-3.5-turbo to rewrite adversarial text as a defense for

*Corresponding author.

LLM-based recommendation models. **Second**, and more fundamentally, current defenses are almost exclusively designed for poisoning-based attacks, where adversaries compromise collaborative signals by injecting fake user profiles or interactions during the training phase. For example, APT [Wu *et al.*, 2021] simulates the poisoning process by injecting fake user data to foster a more robust system. However, in MRSs, attackers can deliberately manipulate the description or image of a targeted item during the inference phase to influence the recommendation results—an aspect not addressed by existing defense methods. To bridge these critical gaps, we propose **Untargeted Adversarial Training with Multimodal Coordination (UAT-MC)**, a novel adversarial training framework specifically designed to defend against evasion-based promotion attacks in MRSs. However, directly applying conventional adversarial training to this setting is non-trivial due to the following two challenges:

C1: Dynamic Attack Target. Unlike poisoning attacks where malicious targets are explicitly injected during training, evasion attacks dynamically select targeted items at the inference phase. This fundamental difference renders traditional targeted adversarial training approaches ineffective, as they rely on pre-defined attack targets to generate adversarial examples, which are unknown during the training phase.

C2: Cross-modal Gradient Mismatch. Under the multi-user promotion setting, MRSs present a unique challenge: combining perturbations across visual and textual modalities often yields suboptimal adversarial examples which degrade the robustness achieved through adversarial training. Cross-modal gradient mismatch arises when visual and textual perturbations in multi-user promotion attacks are dominated by different user groups, causing the two modalities to be optimized toward inconsistent objectives and resulting in misaligned gradient directions.

To tackle the challenge **C1**, we propose **untargeted adversarial training** for MRS, which treats all items as potential targets of evasion-based attacks. Our approach frames recommendation as a multi-label classification task over users, where items serve as labels. Untargeted adversarial training indirectly defends against targeted attacks by globally enhancing the robustness of decision boundaries and disrupting the local gradient information on which attacks rely. The core principle is that the model learns to resist perturbations from any direction, thereby covering attacks originating from specific directions. To tackle challenge **C2**, we propose a novel **gradient-aligned multimodal perturbation method** to address cross-modal gradient mismatch in adversarial training. Unlike perturbing visual and textual features independently, which can diminish attack strength due to misaligned gradients—our method enforces gradient synchronization across modalities by minimizing the cosine distance between visual and textual gradients through a joint loss term. This ensures perturbations coherently push items toward adversarial regions. As shown in later experiments, this coordinated approach maximizes attack potency during training and enhances adversarial robustness against evasion-based promotion attacks.

Our contributions can be summarized as follows:

- We identify evasion-based promotion attacks as a criti-

cal threat to multimodal recommender systems, in which adversaries perturb both the visual and textual features of targeted items.

- We propose a novel Untargeted Adversarial Training with Multimodal Coordination (UAT-MC) framework to enhance the robustness of MRSs against evasion-based promotion attacks. Specifically, the proposed untargeted adversarial training inherently defends against targeted attacks by universally hardening the decision boundaries, thereby covering all potential attack directions.
- We propose gradient-aligned multimodal perturbations to resolve cross-modal gradient mismatch in adversarial training. By minimizing cosine distance between modalities via a joint loss term, our method synchronizes perturbations to maximize attack potency.
- Extensive experiments conducted on real datasets verify the effectiveness of our method.

2 Preliminary

In this section, we conceptually define MRS and describe the evasion-based promotion attacks.

2.1 MRS

Let $\mathcal{U}$ and $\mathcal{I}$ denote user set and item set, respectively. $R \in \{0, 1\}^{|\mathcal{U}| \times |\mathcal{I}|}$ is user-item interaction matrix, where an element $r_{u,i}$ is 1 if there exists interaction between user $u \in \mathcal{U}$ and item $i \in \mathcal{I}$, otherwise 0. Each item i is associated with a visual modality embedding $\mathbf{v}_i$ and a textual modality embedding $\mathbf{t}_i$ which are generated by CNN and Transformer, respectively.

Generally, an MRS f infers the probability $\hat{r}_{u,i}$ that user u will interact with item i based on given R and the modality embeddings $\mathbf{v}_i$ and $\mathbf{t}_i$, i.e.,

$$\hat{r}_{u,i} = f(R, u, i, \mathbf{v}_i, \mathbf{t}_i). \quad (1)$$

Let $\mathcal{D} = \{(u, i_+, i_-)\}$ be a training set, where i_+ is u 's positive sample ($r_{u,i_+} = 1$) and i_- is u 's negative sample ($r_{u,i_-} = 0$). Like traditional recommender systems, an MRS is usually trained with Bayesian Personalized Ranking (BPR) [Rendle *et al.*, 2009] loss function,

$$\mathcal{L}_{\text{BPR}}(\Theta) = \sum_{(u, i_+, i_-) \in \mathcal{D}} -\ln \sigma(\hat{r}_{u,i_+} - \hat{r}_{u,i_-}), \quad (2)$$

where Θ represents the model parameters and $\sigma(\cdot)$ is the sigmoid function.

2.2 Threat Model

Attacker's Objective

Given a well-trained MRS, an evasion-based promotion attack is to imperceptibly perturb the multimodal embeddings of a targeted item i , so that i will appear in the top- K recommendation lists of as many users as possible. Inspired by [Wang *et al.*, 2024b], we define the promotion utility function as:

$$\mathcal{L}_{\text{promotion}} = \frac{1}{N_p} \sum_{u \in \mathcal{U}_p} \text{Sigmoid}(y_{u,i} - y_{u,K}), \quad (3)$$

where $y_{u,K}$ denotes the score of the K -th ranked item for user u . By maximizing Equation (3), the attacker encourages the targeted item score $y_{u,i}$ to surpass the top- K threshold $y_{u,K}$.

Attacker’s Knowledge

In this paper, we assume a white-box scenario for the generation of the adversarial examples in our UAT-MC, where the attacker has full access to the MRS f , including its parameters and gradients. This is because UAT-MC aims to train a more robust MRS, rather than to conduct attacks from the perspective of malicious merchants. This setting represents a worst-case attacker and allows us to evaluate the upper bound of the potential impact of evasion-based promotion attacks.

Attacker’s Capability

Following the ideas of the related works [Tang *et al.*, 2020; Zhang *et al.*, 2021; Guo *et al.*, 2024; Liu *et al.*, 2024; Ong and Khong, 2025], we perturb the multimodal embeddings instead of the raw inputs, which allows more fine-grained and efficient manipulations. Formally, the perturbed embeddings $\mathbf{v}'_i = \mathbf{v}_i + \Delta_v^i$ and $\mathbf{t}'_i = \mathbf{t}_i + \Delta_t^i$. For imperceptibility, the perturbations Δ_m^i ($m \in \{v, t\}$) are constrained by $\|\Delta_m^i\| \leq \epsilon_m$, where ϵ_m is perturbation budget.

3 Cross-modal Gradient Mismatch Analysis

In this section, we demonstrate that under multi-user promotion settings, the gradients of visual and textual perturbations are inherently driven by distinct user groups due to varying modal sensitivities. Consequently, simply aggregating these gradients leads to conflicting optimization directions across modalities, which we term **cross-modal gradient mismatch**. This misalignment fundamentally limits the synergy of multimodal attacks and motivates our UAT-MC framework. We next provide a formal analysis to characterize and verify the existence of this mismatch.

3.1 Objective Inconsistency Across Modalities

Given the promotion objective $\mathcal{L}_{\text{promotion}}$ defined over the user set $\mathcal{U}_p$ (Equation (3)), the optimization problem can be formulated as:

$$\max_{\Delta_v^i, \Delta_t^i} \sum_{u \in \mathcal{U}_p} \mathcal{L}_u(\Delta_v^i, \Delta_t^i), \quad (4)$$

where $\mathcal{L}_u$ denotes the promotion loss contributed by user u .

Accordingly, the perturbations are updated by aggregating gradients from all selected users:

$$G^v = \sum_{u \in \mathcal{U}_p} g_u^v, \quad G^t = \sum_{u \in \mathcal{U}_p} g_u^t, \quad (5)$$

where $g_u^v = \frac{\partial \mathcal{L}_u}{\partial \Delta_v^i}$ and $g_u^t = \frac{\partial \mathcal{L}_u}{\partial \Delta_t^i}$ denote the visual and textual gradients induced by user u , respectively.

To quantify how much each user contributes to the final update direction, we define the **directional contribution** of user u on modality $m \in \{v, t\}$ as:

$$c_u^m = \cos(g_u^m, G^m) \cdot \frac{\|g_u^m\|}{\sum_{u' \in \mathcal{U}_p} \|g_{u'}^m\|}, \quad (6)$$

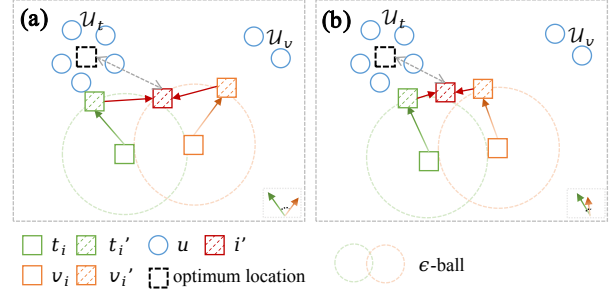


Figure 2: Illustration of objective inconsistency across modalities. (a) Visual and textual perturbations are dominated by user groups $\mathcal{U}_v$ and $\mathcal{U}_t$, leading to conflicting promotion directions after multimodal fusion. (b) With alignment loss, the perturbations from different modalities are constrained to align in a consistent direction, enabling the fused embedding to effectively reach the optimum location.

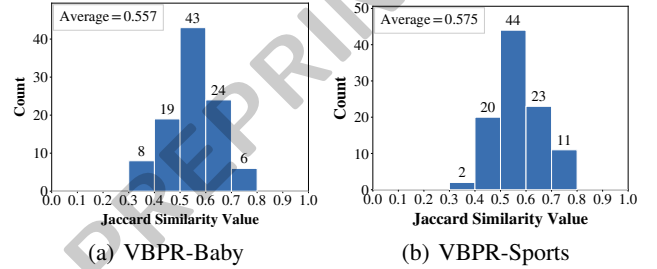


Figure 3: Distribution of Jaccard Similarity between $\mathcal{U}_v$ and $\mathcal{U}_t$.

where $\cos(g_u^m, G^m)$ measures whether user u 's gradient direction is aligned with the aggregated update direction, and $\|g_u^m\|$ reflects the contribution of user u to modality m .

Based on this metric, we define the user set $\mathcal{U}_v$ and the user set $\mathcal{U}_t$ as the top- K users ranked by c_u^v and c_u^t , respectively. To measure the overlap between these two sets, we compute their Jaccard similarity:

$$J(\mathcal{U}_v, \mathcal{U}_t) = \frac{|\mathcal{U}_v \cap \mathcal{U}_t|}{|\mathcal{U}_v \cup \mathcal{U}_t|}. \quad (7)$$

A low Jaccard similarity indicates that the users dominating the visual and textual perturbation updates are largely different, suggesting that the two modalities are implicitly driven by different user groups and, consequently, toward different effective optimization objectives. As illustrated in Fig. 2(a), the visual perturbation Δ_v^i is primarily driven by users in $\mathcal{U}_v$, while the textual perturbation Δ_t^i is dominated by users in $\mathcal{U}_t$. When such modality-specific perturbations are subsequently fused by the MRS, the resulting item representation fails to simultaneously benefit all users, leading to suboptimal promotion outcomes.

3.2 Empirical Evidence of Cross-modal Gradient Mismatch

As mentioned in Section 1, we conduct vanilla FGSM-based promotion attacks on the VBPR model across three datasets. For each dataset, we randomly sample 100 unpopular items

whose interaction counts are no greater than 5. For each targeted item, we identify the user sets $\mathcal{U}_v$ and $\mathcal{U}_t$ using the directional contribution metric in Equation (6), and conduct the promotion attack variants by constructing the promotion loss over $\mathcal{U}_p$, $\mathcal{U}_v$ and $\mathcal{U}_t$, respectively.

The results indicate that aligning the attack objective with modality-specific user groups is crucial, as visual and textual perturbations are driven by distinct user subsets. Specifically, Figure 3 reveals a consistently low Jaccard similarity between $\mathcal{U}_v$ and $\mathcal{U}_t$, confirming their limited overlap. This distinctness is further corroborated by Figure 1, which demonstrates that optimizing over $\mathcal{U}_v$ or $\mathcal{U}_t$ yields significantly higher attack gains than using $\mathcal{U}_p$, as aggregating heterogeneous groups tends to dilute modality-specific optimization signals. (See Appendix B.4 for a detailed case study visualizing this phenomenon on a specific item.)

In summary, cross-modal gradient mismatch limits the effectiveness of multimodal perturbations. Since effective defense requires training against worst-case attacks, it is critical to harmonize optimization directions to maximize attack potency. This motivates our UAT-MC framework, which uses gradient alignment to generate stronger adversarial examples, thereby driving the model to learn more robust representations.

4 Methodology

Figure 4 shows the overview of UAT-MC. In Figure 4, the targeted MRS is divided into two parts, i.e., $f = g \circ h$, where h is the encoder for collaborative filtering and g is the decoder. Usually, h is implemented as a GNN [Wei *et al.*, 2019] or an MLP [He and McAuley, 2016b] for fusion of multimodal embeddings, while g is implemented as an inner product. A training instance consists of user ID embedding $\mathbf{e}_u$, a positive item’s ID embedding $\mathbf{e}_+$, a negative item’s ID embedding $\mathbf{e}_-$, the positive item’s modality embeddings $\mathbf{v}_+$, $\mathbf{t}_+$, and the negative item’s modality embeddings $\mathbf{v}_-$ and $\mathbf{t}_-$.

During the k -th iteration, UAT-MC first generates the adversarial modality embeddings $\mathbf{v}'_+$, $\mathbf{t}'_+$, $\mathbf{v}'_-$, and $\mathbf{t}'_-$, by adding the perturbations generated in the last step. Then through the collaborative filtering encoder h , UAT-MC produces the user embedding $\mathbf{h}_u$ by fusing with the multimodal content of the items interacted with u , the positive and negative item embeddings $\mathbf{h}_+$ and $\mathbf{h}_-$ by fusing their respective multimodal embeddings, and their respective adversarial sample embeddings $\mathbf{h}'_u$, $\mathbf{h}'_+$ and $\mathbf{h}'_-$. Based on these embeddings, the decoder g computes the ranking losses $\mathcal{L}_{\text{BPR}}$ and $\mathcal{L}'_{\text{BPR}}$ on the clean embeddings and the adversarial embeddings, respectively. On these losses, UAT-MC conducts a min-max game to improve the robustness of the MRS f .

It is worth noting that to achieve worst-case adversarial robustness for f , UAT-MC promotes consistent perturbation directions across different modalities by maximizing the gradient alignment loss $\mathcal{L}_{\text{Align}}$ when generating adversarial perturbations. This approach realizes the multimodal coordination, thereby enhancing the aggressiveness of adversarial examples.

4.1 Untargeted Adversarial Training

As mentioned in Section 1, in evasion attacks, the specific targeted item at inference time is unknown during the training phase. To address this challenge, we propose untargeted adversarial training for MRSs, which treats all items as potential targets of evasion-based promotion attacks. This approach is motivated by the fact that recommendation can be formulated as a multi-label classification task, where each item serves as a distinct label. Untargeted adversarial training improves the robustness of the model by strengthening the decision boundary against perturbations from arbitrary directions, thereby effectively defending against perturbations from any direction.

Specifically, given a user u , the attacker adds perturbations $[\Delta_{v_+}, \Delta_{t_+}]$ to the positive item’s multimodal embeddings $[\mathbf{v}_+, \mathbf{t}_+]$ and the perturbations $[\Delta_{v_-}, \Delta_{t_-}]$ to the negative item’s multimodal embeddings $[\mathbf{v}_-, \mathbf{t}_-]$, which result in the adversarial modality embeddings $\mathbf{v}'_+$, $\mathbf{t}'_+$, $\mathbf{v}'_-$, and $\mathbf{t}'_-$. Then the final adversarial embeddings are computed as:

$$\mathbf{h}'_u, \mathbf{h}'_+, \mathbf{h}'_- = h(R, \mathbf{e}_u, \mathbf{e}_+, \mathbf{e}_-, \mathbf{v}'_+, \mathbf{v}'_-, \mathbf{t}'_+, \mathbf{t}'_-). \quad (8)$$

Then the untargeted adversarial training is defined as:

$$\min_{\Theta} \max_{\Delta_v, \Delta_t} \mathcal{L}'_{\text{BPR}}(\Theta, \Delta_v, \Delta_t), \quad (9)$$

where $\mathcal{L}'_{\text{BPR}}(\Theta, \Delta_v, \Delta_t)$ is defined as:

$$\mathcal{L}'_{\text{BPR}}(\Theta, \Delta_v, \Delta_t) = \sum_{(u, i_+, i_-) \in \mathcal{D}} -\ln \sigma(\mathbf{h}'_u \cdot \mathbf{h}'_+ - \mathbf{h}'_u \cdot \mathbf{h}'_-). \quad (10)$$

4.2 Multimodal Coordination

As analyzed in Section 3, independently perturbing each modality in multi-user promotion settings often results in cross-modal gradient mismatch, reducing the effectiveness of adversarial attacks. While identifying and optimizing over specific user subsets (i.e., $\mathcal{U}_v$ and $\mathcal{U}_t$) can improve attack performance, it requires computing user-wise gradients for the entire user base, incurring a prohibitively high computational cost. To circumvent this bottleneck, we propose a lightweight **multimodal coordination** mechanism. Instead of explicitly partitioning users, we focus on directly rectifying the divergent optimization directions in the gradient space. By enforcing gradient-level alignment between modalities, we effectively mitigate the mismatch without the need for costly user identification. This approach achieves coordination with negligible computational overhead, requiring only minimal additional operations during backpropagation.

We first compute the gradients of the adversarial loss with respect to the perturbations in each modality:

$$\Gamma_v = \nabla_{\Delta_v} \mathcal{L}'_{\text{BPR}}, \Gamma_t = \nabla_{\Delta_t} \mathcal{L}'_{\text{BPR}}. \quad (11)$$

Then, we define the alignment loss as the sum of the cosine similarities between the gradients of visual and textual modalities for both positive and negative items:

$$\mathcal{L}_{\text{Align}} = \cos(\Gamma_{v_+}, \Gamma_{t_+}) + \cos(\Gamma_{v_-}, \Gamma_{t_-}). \quad (12)$$

Maximizing $\mathcal{L}_{\text{Align}}$ encourages the perturbations in different modalities to be directionally aligned, thereby pushing the targeted item toward the adversarial region.

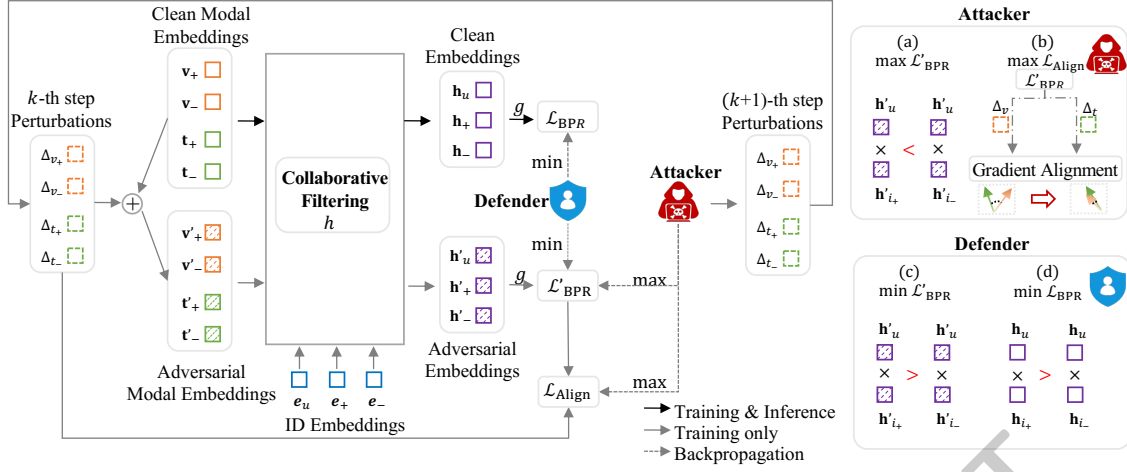


Figure 4: The framework of UAT-MC.

Algorithm 1 UAT-MC

Input: Training data $\mathcal{D}$; Learning rate η ; Perturbation budgets ϵ_t, ϵ_v ; Hyperparameters λ, α, β
Output: Robust MRS model parameters Θ

- 1: Initialize Θ from normal-trained MRS
- 2: **while** not converged **do**
- 3: Randomly draw an example (u, i_+, i_-) from $\mathcal{D}$
- 4: // Max phase
- 5: $\mathcal{L}'_{\text{BPR}} = -\ln \sigma(\mathbf{h}'_u \cdot \mathbf{h}'_+ - \mathbf{h}'_u \cdot \mathbf{h}'_-)$
- 6: $\Gamma_v = \nabla_{\Delta_v} \mathcal{L}'_{\text{BPR}}, \Gamma_t = \nabla_{\Delta_t} \mathcal{L}'_{\text{BPR}}$
- 7: $\mathcal{L}_{\text{Align}} = \cos(\Gamma_{v_+}, \Gamma_{t_+}) + \cos(\Gamma_{v_-}, \Gamma_{t_-})$
- 8: $\mathcal{L}_{\text{max}} = \mathcal{L}'_{\text{BPR}} + \alpha \mathcal{L}_{\text{Align}}$
- 9: $\Delta_v \leftarrow \epsilon_v \cdot \frac{\nabla_{\Delta_v} \mathcal{L}'_{\text{max}}}{\|\nabla_{\Delta_v} \mathcal{L}'_{\text{max}}\|_2}, \Delta_t \leftarrow \epsilon_t \cdot \frac{\nabla_{\Delta_t} \mathcal{L}'_{\text{max}}}{\|\nabla_{\Delta_t} \mathcal{L}'_{\text{max}}\|_2}$
- 10: // Min phase
- 11: $\mathcal{L}_{\text{min}} = \mathcal{L}_{\text{BPR}} + \lambda \mathcal{L}'_{\text{BPR}}(\Delta_v, \Delta_t) + \beta \|\Theta\|_2$
- 12: $\Theta \leftarrow \Theta - \eta \nabla_{\Theta} \mathcal{L}_{\text{min}}$
- 13: **end while**
- 14: **return** Θ

Dataset	#Users	#Items	#Interactions	Sparsity
Baby	19,445	7,037	160,792	99.883%
Sports	35,598	18,357	296,337	99.955%
Clothing	39,387	23,033	278,677	99.969%

Table 1: Statistics of the three experimental datasets.

4.3 Optimization

The training is divided into two phases:

- **Attacking** (Max phase): Generate the most disruptive perturbations Δ_v and Δ_t within ℓ_2 -norm budgets ϵ_v and ϵ_t by maximizing $\lambda \mathcal{L}'_{\text{BPR}} + \alpha \mathcal{L}_{\text{Align}}$.
- **Defending** (Min phase): Update model parameters Θ by minimizing $\mathcal{L}_{\text{min}} = \mathcal{L}_{\text{BPR}} + \lambda \mathcal{L}'_{\text{BPR}}(\Delta_v, \Delta_t) + \beta \|\Theta\|_2$.

The overall training process is described in Algorithm 1, where the MRS is pre-trained using the standard BPR loss.

5 Experiments

The objectives of experiments are to answer the following research questions:

- **RQ1:** Can UAT-MC defend against evasion-based promotion attacks?
- **RQ2:** Does multimodal coordination enhance the effectiveness of adversarial defense?
- **RQ3:** How does the perturbation budget in adversarial training affect the defense performance?
- **RQ4:** How do the hyper-parameters affect the trade-off between recommendation performance and adversarial robustness?

All experiments are implemented using PyTorch 2.2.0 and run on an NVIDIA RTX 4090 GPU with 48GB of memory.

5.1 Experimental Settings**Datasets**

We conduct experiments on three Amazon datasets Baby, Sports and Clothing [He and McAuley, 2016a], which have recently been adopted in MRSs research [Yu *et al.*, 2023; Guo *et al.*, 2024; Li *et al.*, 2025]. The dataset statistics are summarized in Table 1. In our work, we use the pre-extracted visual features and textual features that have been published in [Zhou *et al.*, 2023c; Zhou *et al.*, 2023a].

Evaluation Metrics

To measure recommendation performance before and after adversarial training, we use Recall@10 and NDCG@10 under the leave-one-out evaluation protocol following standard evaluation settings in prior works [Zhou *et al.*, 2023c; Guo *et al.*, 2024]. To evaluate the effectiveness of the promotion attack, we adopt the hit rate $\text{Hit}_i@K = N_{\text{rec}}^i / |\mathcal{U}| \cdot 100\%$ and the relative improvement ratio $\text{Gain}_{\text{Hit}_i@K} = (\text{Hit}_{\text{after}}@K - \text{Hit}_{\text{before}}@K) / \text{Hit}_{\text{before}}@K \cdot 100\%$ as evaluation metrics, where N_{rec}^i is the number of users for whom the targeted item i appears in the top- K recommendation list (with K defaulted to 50), $\text{Hit}_{\text{after}}@K$ and $\text{Hit}_{\text{before}}@K$ denote the hit rate of the targeted item after and before the attack, respectively.

Dataset	Victim Model	Defense Method	Hit _{before}	FGSM-based Attack				PGD-based Attack				Recommendation	
				$\mathcal{L}_{\text{prom}}$		$\mathcal{L}_{\text{prom}} + \mathcal{L}_{\text{Align}}$		$\mathcal{L}_{\text{prom}}$		$\mathcal{L}_{\text{prom}} + \mathcal{L}_{\text{Align}}$		Recall	NDCG
				Hit _{after}	Gain _{Hit}	Hit _{after}	Gain _{Hit}	Hit _{after}	Gain _{Hit}	Hit _{after}	Gain _{Hit}		
Baby	VBPR	w/o AT	0.667%	3.865%	479.38%	3.908%	485.80%	4.111%	516.15%	4.143%	520.99%	0.0503	0.027
		UAT	0.622%	2.381%	282.96%	2.383%	283.41%	2.450%	294.15%	2.451%	294.25%	0.0484	0.0264
		UAT-MC	0.621%	1.515%	143.98%	1.515%	144.00%	1.544%	148.65%	1.544%	148.65%	0.0476	0.0253
	MMGCN	w/o AT	0.622%	1.747%	181.03%	1.748%	181.14%	3.399%	446.76%	3.421%	450.30%	0.0413	0.022
		UAT	0.472%	0.507%	7.29%	0.508%	7.49%	0.515%	9.12%	0.516%	9.34%	0.0344	0.0185
		UAT-MC	0.479%	0.513%	7.21%	0.514%	7.47%	0.519%	8.46%	0.520%	8.62%	0.0341	0.0185
Sports	VBPR	w/o AT	0.025%	6.547%	25722.99%	6.586%	25876.73%	7.253%	28506.93%	7.279%	28610.25%	0.0586	0.0319
		UAT	0.025%	0.028%	13.45%	0.028%	13.45%	0.029%	13.73%	0.029%	13.73%	0.0510	0.0280
		UAT-MC	0.024%	0.026%	10.32%	0.026%	10.32%	0.026%	10.32%	0.026%	10.32%	0.0508	0.0280
	MMGCN	w/o AT	0.024%	0.266%	1024.04%	0.268%	1033.23%	0.572%	2315.13%	0.581%	2352.82%	0.0385	0.0206
		UAT	0.035%	0.038%	7.01%	0.039%	11.02%	0.044%	24.45%	0.044%	25.45%	0.0313	0.0172
		UAT-MC	0.034%	0.034%	0.00%	0.037%	8.44%	0.041%	19.14%	0.041%	19.14%	0.0317	0.0173
Clothing	VBPR	w/o AT	0.036%	1.335%	3558.26%	1.380%	3682.09%	1.419%	3788.87%	1.438%	3838.61%	0.0384	0.0211
		UAT	0.031%	0.287%	818.66%	0.303%	869.17%	0.306%	878.50%	0.312%	898.38%	0.0343	0.0189
		UAT-MC	0.030%	0.038%	26.64%	0.038%	26.85%	0.038%	27.70%	0.038%	27.70%	0.0333	0.0184
	MMGCN	w/o AT	0.014%	1.886%	13282.43%	1.941%	13670.72%	9.035%	64019.37%	9.134%	64715.32%	0.0239	0.0123
		UAT	0.014%	0.025%	76.00%	0.026%	81.78%	0.035%	148.44%	0.036%	149.33%	0.0221	0.0114
		UAT-MC	0.015%	0.019%	29.44%	0.020%	33.77%	0.032%	119.05%	0.032%	121.65%	0.0221	0.0115

Table 2: Performance of promotion attacks on two classical models. All reported numbers are averaged results. The best defense results are highlighted in bold. (Note: Gain_{Hit@50} values are computed using unrounded Hit@50 scores).

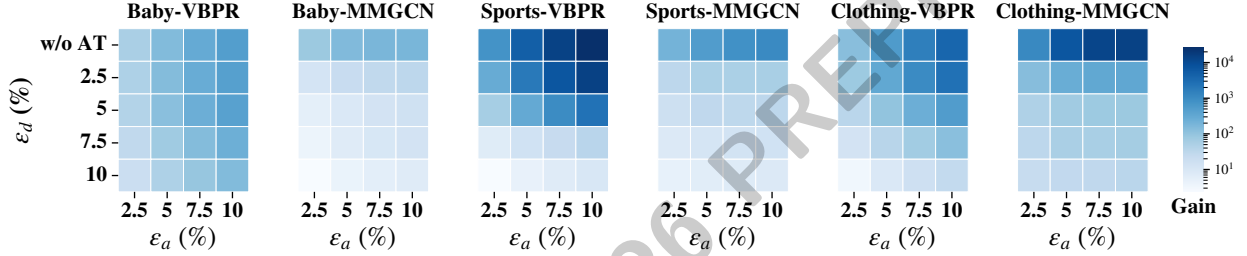


Figure 5: Visualization of attack effectiveness under varying budgets. The heatmaps display the $\text{Gain}_{\text{Hit@50}}$ (log scale) across three datasets and two models. The x-axis represents the attack budget ϵ_a , and the y-axis represents the defense budget ϵ_d . Darker colors indicate higher attack gains (weaker defense), while lighter colors indicate effective defense.

A larger $\text{Gain}_{\text{Hit@50}}$ indicates a more effective promotion attack, suggesting weaker model robustness.

Targeted Items

At the MRS’s inference phase, we randomly select 100 targeted items from unpopular items, whose interaction count is 5. We conduct promotion attacks on each targeted item individually and record the average $\text{Hit}_i@K$ and $\text{Gain}_{\text{Hit}_i@K}$.

Victim Models

We choose two mainstream and classical multimodal recommendation models as our victim models:

- VBPR [He and McAuley, 2016b]: As a representative of MLP-based models, VBPR incorporates visual features for user preference learning with BPR loss. Following [Zhou *et al.*, 2023c; Guo *et al.*, 2024], we concatenate the multimodal embeddings and the item ID embedding as the item embeddings.
- MMGCN [Wei *et al.*, 2019]: As a representative of GCN-based models, MMGCN constructs three modal-specific graph to learn different modality features. It concatenates all modality embeddings to obtain the representations of users or items.

We follow MMRec [Zhou *et al.*, 2023a] to save the model parameters at the point of best performance. The hyper-parameters search spaces are provided in the Appendix B.1.

Promotion Attacks

We conduct promotion attacks by maximizing $\mathcal{L}_{\text{promotion}}$ (Eq. 3) using two standard strategies: FGSM [Goodfellow *et al.*, 2014] and PGD [Madry *et al.*, 2017]. FGSM serves as a single-step baseline, whereas PGD functions as a stronger iterative attacker. In our experiments, the PGD attack is configured with 10 iterations and a step size of $\alpha = 1.25 \cdot \epsilon_m / 10$.

Defenders

To verify the effectiveness of the Untargeted Adversarial Training and the Multimodal Coordination, we compare UAT-MC with two baseline models: one is the victim model without any defense mechanisms, denoted as **w/o AT**, and the other is the variant of UAT-MC, denoted as **UAT**, which applies perturbations to both visual and textual modalities independently without coordination.

Implementation Details

During adversarial training phase, we set the maximum perturbation magnitude ϵ_m as 10% of the 2-norm of the in-

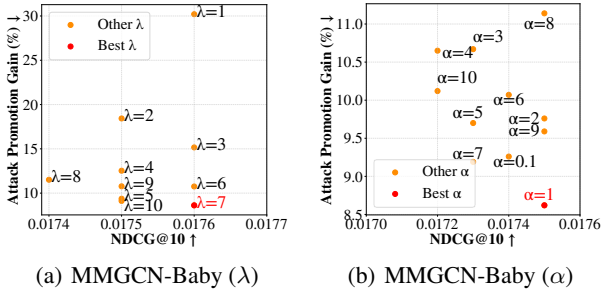


Figure 6: Trade-off between recommendation performance (NDCG@10) and adversarial robustness (GainHit@50) under PGD(w/ $\mathcal{L}_{\text{Align}}$)-based attack with varying λ and α on the Baby dataset.

put embedding for modality m , the coefficient α is searched in $\{0.1, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10\}$, the coefficient λ is searched in $\{1, 2, 3, 4, 5, 6, 7, 8, 9, 10\}$.

5.2 Defense Effect (RQ1 and RQ2)

We conduct a systematic evaluation of the performance of different defense methods on two attack scenarios (FGSM-based and PGD-based) on two mainstream MRSs (VBPR [He and McAuley, 2016b] and MMGCN [Wei *et al.*, 2019]). The victim models take the original and perturbed multimodal embeddings as input, respectively. Meanwhile, we also report the recommendation performance on the clean test set to reflect the impact of adversarial training on overall recommendation quality. The results are shown in Table 2, from which we have the following observations:

- **Across all dataset-model combinations, unpopular items consistently exhibit vulnerability to promotion attacks (w/o $\mathcal{L}_{\text{Align}}$),** which risk misleading the target model into treating them as popular items. It can be observed that perturbations generated by PGD are more effective in promoting the targeted items compared to those generated by FGSM. Statistically, on the Baby dataset, the average Hit@50 of unpopular items increases from 0.667% to 3.865% under FGSM-based perturbations, and further to 4.111% under PGD-based perturbations. This implies that more than 3.444% of users will be misled into recommending an item (e.g., targeted item) that should not appear in the recommendation lists.
- **Attacks equipped with $\mathcal{L}_{\text{Align}}$ consistently achieve higher promotion gains than their counterparts without alignment,** indicating that explicitly coordinating visual and textual perturbations leads to more effective promotion. This improvement is particularly evident for stronger attacks. For example, under PGD-based attacks on the Baby dataset, introducing $\mathcal{L}_{\text{Align}}$ increases the GainHit@50 of VBPR from 516.15% to 520.99%, and that of MMGCN from 446.76% to 450.30%. Similar trends are observed on the Sports and Clothing datasets, where alignment consistently yields additional promotion gains across both VBPR and MMGCN.
- **After untargeted adversarial training (UAT), the MRSs exhibit significantly enhanced resistance to**

evasion-based promotion attacks. The results show that the effectiveness of both FGSM-based and PGD-based attacks drops significantly after UAT. In the following, we analyze based on the average GainHit@50 results of the two attacks (w/ $\mathcal{L}_{\text{Align}}$). On the Baby dataset, UAT reduces the GainHit@50 of VBPR from 503.39% to 288.83%, achieving a 42.62% relative reduction. For MMGCN, the value drops significantly from 315.72% to 8.42%, corresponding to a 97.33% reduction. Similar trends are observed on the Sports dataset, where VBPR and MMGCN experience reductions of 99.95% and 98.92%, respectively. On the Clothing dataset, UAT reduces GainHit@50 by 76.50% for VBPR and 99.71% for MMGCN.

- **UAT with Multimodal Coordination can further enhance its defense capability.** According to our observations, the GainHit@50 of MMGCN on the Sports dataset drops from 25.45% to 19.14%, while that of VBPR on the Clothing dataset decreases significantly from 898.38% to 27.70% after applying multimodal coordination. Notably, the improvement in robustness is achieved with minimal impact on recommendation performance.

5.3 Hyper-parameter Study (RQ3 and RQ4)

This section explores how the perturbation budget ϵ_m and the hyper-parameters λ and α in adversarial training influence the trade-off between adversarial robustness and recommendation performance.

- **Impact of ϵ_d and ϵ_a** To examine the robustness and generalization ability of UAT-MC, we conduct adversarial training with different defense perturbation budgets $\epsilon_d \in \{2.5\%, 5\%, 7.5\%, 10\%\}$, and evaluate the trained models under FGSM-based promotion attacks with varying attack budgets $\epsilon_a \in \{2.5\%, 5\%, 7.5\%, 10\%\}$. The results are visualized in Figure 5, where specific numerical values are detailed in the Appendix B.3. We observe a clear transition from dark blue to light blue/white as ϵ_d increases, indicating that our defense method effectively mitigates the promotion attack.
- **Impact of λ and α** We investigate the trade-off between adversarial robustness and recommendation performance by varying $\lambda \in [1, 10]$ and $\alpha \in [0.1, 10]$. Figure 6 (Baby dataset) and the Appendix B.1 (others) visualize this relationship, plotting recommendation performance (NDCG@10) against attack gain (GainHit@50). A clear trade-off is observed: higher recommendation performance typically comes at the cost of increased vulnerability.

6 Conclusion

In this work, we address the vulnerability of MRSs to promotion attacks. Crucially, we verify the existence of cross-modal gradient mismatch in multi-user promotion settings and proposed UAT-MC to mitigate it via a novel gradient alignment regularization. Extensive experiments demonstrate the effectiveness of UAT-MC in defending against evasion-based promotion attacks.

Acknowledgments

This work is supported by Natural Science Foundation of Sichuan Province under grant 2024NSFSC0516 and National Natural Science Foundation of China under grant 61972270.

References

- [Chen *et al.*, 2024] Lijian Chen, Wei Yuan, Tong Chen, Guanhua Ye, Nguyen Quoc Viet Hung, and Hongzhi Yin. Adversarial item promotion on visually-aware recommender systems by guided diffusion. *ACM Trans. Inf. Syst.*, 42(6), August 2024.
- [Di Noia *et al.*, 2020] Tommaso Di Noia, Daniele Malitesta, and Felice Antonio Merra. Taamr: Targeted adversarial attack against multimedia recommender systems. In *2020 50th Annual IEEE/IFIP International Conference on Dependable Systems and Networks Workshops (DSN-W)*, pages 1–8, 2020.
- [Goodfellow *et al.*, 2014] Ian J Goodfellow, Jonathon Shlens, and Christian Szegedy. Explaining and harnessing adversarial examples. *arXiv preprint arXiv:1412.6572*, 2014.
- [Guo *et al.*, 2024] Zhiqiang Guo, Jianjun Li, Guohui Li, Chaoyang Wang, Si Shi, and Bin Ruan. Lgmrec: local and global graph learning for multimodal recommendation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 8454–8462, 2024.
- [He and McAuley, 2016a] Ruining He and Julian McAuley. Ups and downs: Modeling the visual evolution of fashion trends with one-class collaborative filtering. In *proceedings of the 25th international conference on world wide web*, pages 507–517, 2016.
- [He and McAuley, 2016b] Ruining He and Julian McAuley. Vbpr: visual bayesian personalized ranking from implicit feedback. In *Proceedings of the AAAI conference on artificial intelligence*, volume 30, 2016.
- [Hsiao *et al.*, 2022] Shao-Ping Hsiao, Yu-Che Tsai, and Cheng-Te Li. Unsupervised post-time fake social message detection with recommendation-aware representation learning. In *Companion Proceedings of the Web Conference 2022*, pages 232–235, 2022.
- [Li *et al.*, 2025] Hongji Li, Hanwen Du, Youhua Li, Junchen Fu, Chunxiao Li, Ziyi Zhuang, Jiakang Li, and Yongxin Ni. Teach me how to denoise: A universal framework for denoising multi-modal recommender systems via guided calibration. In *Proceedings of the Eighteenth ACM International Conference on Web Search and Data Mining, WSDM ’25*, page 782–791, New York, NY, USA, 2025. Association for Computing Machinery.
- [Lin *et al.*, 2023] Weilin Lin, Xiangyu Zhao, Yejing Wang, Yuanshao Zhu, and Wanyu Wang. Autodenoise: Automatic data instance denoising for recommendations. In *Proceedings of the ACM Web Conference 2023*, pages 1003–1011, 2023.
- [Liu and Larson, 2021] Zhuoran Liu and Martha Larson. Adversarial item promotion: Vulnerabilities at the core of top-n recommenders that use images to address cold start. In *Proceedings of the Web Conference 2021, WWW ’21*, page 3590–3602, New York, NY, USA, 2021. Association for Computing Machinery.
- [Liu *et al.*, 2024] Yifan Liu, Kangning Zhang, Xiangyuan Ren, Yanhua Huang, Jiarui Jin, Yingjie Qin, Ruilong Su, Ruiwen Xu, Yong Yu, and Weinan Zhang. Alignrec: Aligning and training in multimodal recommendations. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management, CIKM ’24*, page 1503–1512, New York, NY, USA, 2024. Association for Computing Machinery.
- [Madry *et al.*, 2017] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. *arXiv preprint arXiv:1706.06083*, 2017.
- [Mu *et al.*, 2025] Lingyu Mu, Zhengxiao Liu, Zhitong Zhu, and Zheng Lin. Trust-grs: A trustworthy training framework for graph neural network based recommender systems against shilling attacks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 12408–12416, 2025.
- [Nguyen Thanh *et al.*, 2023] Toan Nguyen Thanh, Nguyen Duc Khang Quach, Thanh Tam Nguyen, Thanh Trung Huynh, Viet Hung Vu, Phi Le Nguyen, Jun Jo, and Quoc Viet Hung Nguyen. Poisoning gnn-based recommender systems with generative surrogate-based attacks. *ACM Trans. Inf. Syst.*, 41(3), February 2023.
- [Ong and Khong, 2025] Rongqing Kenneth Ong and Andy W. H. Khong. Spectrum-based modality representation fusion graph convolutional network for multimodal recommendation. In *Proceedings of the Eighteenth ACM International Conference on Web Search and Data Mining, WSDM ’25*, page 773–781, New York, NY, USA, 2025. Association for Computing Machinery.
- [Rendle *et al.*, 2009] Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. Bpr: Bayesian personalized ranking from implicit feedback. In *Proceedings of the Twenty-Fifth Conference on Uncertainty in Artificial Intelligence, UAI ’09*, page 452–461, Arlington, Virginia, USA, 2009. AUAI Press.
- [Sun *et al.*, 2020] Rui Sun, Xuezhi Cao, Yan Zhao, Junchen Wan, Kun Zhou, Fuzheng Zhang, Zhongyuan Wang, and Kai Zheng. Multi-modal knowledge graphs for recommender systems. In *Proceedings of the 29th ACM international conference on information & knowledge management*, pages 1405–1414, 2020.
- [Tang *et al.*, 2020] Jinhui Tang, Xiaoyu Du, Xiangnan He, Fajie Yuan, Qi Tian, and Tat-Seng Chua. Adversarial training towards robust multimedia recommender system. *IEEE Transactions on Knowledge and Data Engineering*, 32(5):855–867, 2020.
- [Wang *et al.*, 2022] Shilei Wang, Peng Zhang, Hui Wang, Hongtao Yu, and Fuzhi Zhang. Detecting shilling groups

- in online recommender systems based on graph convolutional network. *Information Processing & Management*, 59(5):103031, 2022.
- [Wang *et al.*, 2023] Qifan Wang, Yinwei Wei, Jianhua Yin, Jianlong Wu, Xuemeng Song, and Liqiang Nie. Dualgnn: Dual graph neural network for multimedia recommendation. *Multimedia, IEEE Trans. on (T-MM)*, 25(000):11, 2023.
- [Wang *et al.*, 2024a] Zongwei Wang, Min Gao, Junliang Yu, Xinyi Gao, Quoc Viet Hung Nguyen, Shazia Sadiq, and Hongzhi Yin. Llm-powered text simulation attack against id-free recommender systems, 2024.
- [Wang *et al.*, 2024b] Zongwei Wang, Junliang Yu, Min Gao, Hongzhi Yin, Bin Cui, and Shazia Sadiq. Unveiling vulnerabilities of contrastive recommender systems to poisoning attacks. In *Proceedings of the 30th ACM SIGKDD conference on knowledge discovery and data mining*, pages 3311–3322, 2024.
- [Wei *et al.*, 2019] Yinwei Wei, Xiang Wang, Liqiang Nie, Xiangnan He, Richang Hong, and Tat-Seng Chua. Mmgn: Multi-modal graph convolution network for personalized recommendation of micro-video. In *Proceedings of the 27th ACM international conference on multimedia*, pages 1437–1445, 2019.
- [Wei *et al.*, 2020] Yinwei Wei, Xiang Wang, Liqiang Nie, Xiangnan He, and Tat-Seng Chua. Graph-refined convolutional network for multimedia recommendation with implicit feedback. In *Proceedings of the 28th ACM International Conference on Multimedia*, MM ’20, page 3541–3549, New York, NY, USA, 2020. Association for Computing Machinery.
- [Wu *et al.*, 2021] Chenwang Wu, Defu Lian, Yong Ge, Zhihao Zhu, Enhong Chen, and Senchao Yuan. Fight fire with fire: Towards robust recommender systems via adversarial poisoning training. In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR ’21, page 1074–1083, New York, NY, USA, 2021. Association for Computing Machinery.
- [Wu *et al.*, 2023] Chenwang Wu, Defu Lian, Yong Ge, Zhihao Zhu, and Enhong Chen. Influence-driven data poisoning for robust recommender systems. *IEEE Trans. Pattern Anal. Mach. Intell.*, 45(10):11915–11931, October 2023.
- [Yang *et al.*, 2024] Shiyi Yang, Chen Wang, Xiwei Xu, Liming Zhu, and Lina Yao. Attacking visually-aware recommender systems with transferable and imperceptible adversarial styles. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management*, CIKM ’24, page 2900–2909, New York, NY, USA, 2024. Association for Computing Machinery.
- [Yu *et al.*, 2023] Penghang Yu, Zhiyi Tan, Guanming Lu, and Bing-Kun Bao. Multi-view graph convolutional network for multimedia recommendation. In *Proceedings of the 31st ACM international conference on multimedia*, pages 6576–6585, 2023.
- [Zhang *et al.*, 2021] Jinghao Zhang, Yanqiao Zhu, Qiang Liu, Shu Wu, Shuhui Wang, and Liang Wang. Mining latent structures for multimedia recommendation. In *Proceedings of the 29th ACM international conference on multimedia*, pages 3872–3880, 2021.
- [Zhang *et al.*, 2024] Jinghao Zhang, Yuting Liu, Qiang Liu, Shu Wu, Guibing Guo, and Liang Wang. Stealthy attack on large language model based recommendation. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar, editors, *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5839–5857, Bangkok, Thailand, August 2024. Association for Computational Linguistics.
- [Zhou *et al.*, 2023a] Hongyu Zhou, Xin Zhou, Zhiwei Zeng, Lingzi Zhang, and Zhiqi Shen. A comprehensive survey on multimodal recommender systems: Taxonomy, evaluation, and future directions, 2023.
- [Zhou *et al.*, 2023b] Hongyu Zhou, Xin Zhou, Lingzi Zhang, and Zhiqi Shen. Enhancing dyadic relations with homogeneous graphs for multimodal recommendation. In *ECAI 2023*, pages 3123–3130. IOS Press, 2023.
- [Zhou *et al.*, 2023c] Xin Zhou, Hongyu Zhou, Yong Liu, Zhiwei Zeng, Chunyan Miao, Pengwei Wang, Yuan You, and Feijun Jiang. Bootstrap latent representations for multi-modal recommendation. In *Proceedings of the ACM web conference 2023*, pages 845–854, 2023.