

Embedding-Space Orthogonal Decomposition for Robust Social Recommendation

Rongfeng Guo
2023096062@email.szu.edu.cn
CSSE, Shenzhen University
Shenzhen, Guangdong, China

Yinxuan Huang
huangyinxuan22@nudt.edu.cn
National University of Defense
Technology
Changsha, Hunan, China

Wei Chen
weic@microsoft.com
Microsoft Research
Beijing, China

Mingyang Zhou
zmy@szu.edu.cn
CSSE, Shenzhen University
Shenzhen, Guangdong, China

Yusen Wu
3231319130@smail.fjut.edu.cn
Fujian University of Technology
Fuzhou, Fujian, China

Yangchen Zeng
zengyangchen@foxmail.com
Southeast University
Nanjing, Jiangsu, China

Han Chen
2023096053@email.szu.edu.cn
Shenzhen University
Shenzhen, Guangdong, China

Hao Liao
haoliao@szu.edu.cn
CSSE, Shenzhen University
Shenzhen, Guangdong, China

Abstract

Graph-based social recommendation leverages both the interaction graph and the social graph to model user preferences, especially under sparse feedback. However, users' intricate social behaviors may introduce mismatched social ties that contaminate user representations and harm the models' robustness. The majority of existing methods mitigate this by pruning, rewiring, or assigning edge-wise weights before social aggregation. From users' historical behaviors, we observe that a social neighbor often overlaps with the target user on specific interests but differs in others. Thus, using a single weight for each social connection is insufficient, as it only scales the overall message intensity and fails to selectively suppress the misaligned components within the aggregated message. To fill this gap, we propose **Orthogonal Decomposition for Social Recommendation (ODSR)**, an embedding-space framework that orthogonally decomposes the aggregated social message into an aligned component and an orthogonal deviation, and learns a dimension-wise vector gate to regulate the deviation under ranking supervision. Additionally, we introduce a contrastive regularizer that perturbs representations along deviation directions to enhance robustness against imperfect social signals. Extensive experiments on three datasets show that ODSR consistently outperforms strong baselines, and additional analyses verify the effectiveness of selectively gating the orthogonal deviation.

CCS Concepts

• **Information systems** → **Social recommendation**; *Recommender systems*; • **Computing methodologies** → *Learning latent representations*; *Neural networks*.

Keywords

Social recommendation, graph neural networks, orthogonal decomposition, representation learning

ACM Reference Format:

Rongfeng Guo, Yinxuan Huang, Wei Chen, Mingyang Zhou, Yusen Wu, Yangchen Zeng, Han Chen, and Hao Liao. 2026. Embedding-Space Orthogonal Decomposition for Robust Social Recommendation. In *Proceedings of the 32nd ACM SIGKDD Conference on Knowledge Discovery and Data Mining V.2 (KDD '26)*, August 09–13, 2026, Jeju Island, Republic of Korea. ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/3770855.3817864>

Resource Availability: Code is available at <https://doi.org/10.5281/zenodo.20424225> and <https://github.com/Rongfeng-Guo/ODSR-KDD-2026>.

1 Introduction

With the proliferation of social platforms, social recommendation has become an increasingly practical paradigm for personalized recommendation, where user preferences are inferred by jointly exploiting the interaction graph and the social graph [20, 21]. By modeling observed feedback as a graph and performing message passing, graph-based recommenders propagate collaborative signals to learn representations that support effective item ranking. Social recommendation further incorporates social relations to enrich user representations, especially in scenarios with sparse feedback, where learning solely from interactions may lack robustness. Motivated by homophily [24] and classic social-influence theories [4, 5], as well as their learning-based instantiations in social recommendation and social networks [32, 34, 37], many recent methods combine social aggregation with self-supervised regularization to stabilize representation learning in sparse regimes.

A key challenge is that social neighbors are often only *partially aligned* with the target user. Figure 1 illustrates a toy example. The target user John shares only part of his neighbors' intents: Tony overlaps with John on *basketball* but introduces off-anchor interests such as *painting* and *cooking*, while Lisa overlaps on *basketball* and *tennis* but brings in *fashion* as noise. This implies that a social



This work is licensed under a Creative Commons Attribution 4.0 International License. *KDD '26, Jeju Island, Republic of Korea*
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2259-2/2026/08
<https://doi.org/10.1145/3770855.3817864>

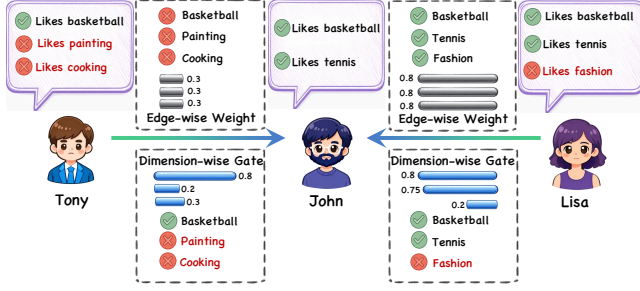


Figure 1: Motivation of ODSR. John has partial alignment with neighbors Tony and Lisa. (Top) Edge-wise weight uniformly scales all dimensions. (Bottom) Dimension-wise gate selectively suppresses noise while preserving aligned signals.

message can be simultaneously helpful and harmful across different intent dimensions. We define *partial alignment* as the phenomenon that a single social neighbor is aligned with the target user on some intent dimensions but deviates on others. This mixture is *within-embedding*: aligned and deviating components coexist within the aggregated message, but in different dimensions.

Robustness in social recommendation is traditionally studied by intervening on social graph edges, such as filtering edges or assigning an edge-wise reliability weight [12, 25], which we refer to as the *edge-space* perspective. However, edge unreliability is not the whole story [19, 53], since partial alignment can hold even for a useful edge when the neighbor is aligned on some embedding dimensions but deviates on others. Alternatively, one can intervene on the aggregated neighbor message by controlling which embedding dimensions are absorbed or suppressed, which we refer to as the *embedding space* perspective.

Methods in the edge-space perspective work well when a neighbor is largely beneficial or largely harmful, because the edge can be treated as uniformly reliable or uniformly unreliable. However, under *partial alignment*, edge-space control becomes coarse because it applies the same decision to all embedding dimensions carried by the neighbor message. Scalar reliability scores or attention weights act as a *uniform multiplier* over all dimensions, so they cannot suppress the deviating dimensions while preserving the aligned dimensions. This leads to the misalignment of signals, even when the edge is beneficial, resulting in suboptimal robustness and ranking accuracy. We quantify *partial alignment* via the *directed intent overlap* $r_{u \leftarrow v}^+ \in [0, 1]$, defined as the fraction of v 's interaction-induced intent groups that are also covered by u (Sec. 2). Mid-range r^+ accounts for the majority of ties across datasets, indicating that most neighbors are neither fully aligned nor fully misaligned, and their aggregated messages therefore contain both aligned signals and off-anchor deviations (Figure 1). In short, the challenge is within-embedding mixture under partial alignment, and our key idea is to control the injected social message in embedding space via orthogonal decomposition and dimension-wise deviation gating.

To address this limitation, **ODSR** (Orthogonal Decomposition for Social Recommendation) is developed as an embedding space framework for within-message mixture under partial alignment, rather than edge-wise scaling in the edge-space perspective. Instead of relying on edge-space scaling, ODSR decomposes the aggregated social message into an aligned component and an orthogonal deviation, then learns a dimension-wise gating vector under ranking supervision to regulate the deviation. Furthermore, a deviation-aware contrastive regularizer perturbs representations along orthogonal deviation directions, improving robustness to imperfect social signals. In summary, our main contributions are as follows:

- We quantify partial alignment via directed intent overlap and show that mid-range overlap dominates across datasets, revealing mixture within a message as a prevalent robustness bottleneck in social aggregation.
- We propose **ODSR**, an embedding-space framework that learns a dimension-wise gating vector to regulate orthogonal deviation under ranking supervision, together with deviation-aware contrastive regularization.
- On three datasets, ODSR yields consistent improvements over strong baselines, and further analyses verify that selectively gating orthogonal deviations is the key driver.

2 Data Analysis

We conduct training-split analyses to motivate why social regulation should operate inside the embedding space rather than only on graph edges. Specifically, we ask two questions. First, are social ties usually fully aligned or only partially aligned with the target user? Second, if a social message contains an off-anchor part, is this part uniformly spread across coordinates or concentrated on a subset of coordinates in the learned basis? The first question is studied by directed intent overlap $r_{u \leftarrow v}^+$, and the second is studied by the coordinate concentration of anchor-relative deviation. Together, the two analyses explain why a scalar edge weight is insufficient and why a dimension-wise gate is useful.

2.1 Partial Alignment in Social Ties

Problem Definition. We study social recommendation with users $\mathcal{U}$, items $\mathcal{V}$, interactions $\mathcal{E}_I \subseteq \mathcal{U} \times \mathcal{V}$, and directed social ties $\mathcal{E}_S \subseteq \mathcal{U} \times \mathcal{U}$, forming an interaction graph $\mathcal{G}_I = (\mathcal{U} \cup \mathcal{V}, \mathcal{E}_I)$ and a social graph $\mathcal{G}_S = (\mathcal{U}, \mathcal{E}_S)$. A directed tie $(u, v) \in \mathcal{E}_S$ is written as $u \leftarrow v$, meaning that neighbor v injects a social message into target user u . Graph recommenders learn embeddings by layer-wise message passing; we denote layer- l user/item embeddings by $\mathbf{h}_u^{(l)}, \mathbf{h}_i^{(l)} \in \mathbb{R}^D$, with D denoting the embedding dimensionality and $l=0$ denoting trainable ID embeddings. The ranking goal is top- K retrieval with dot-product scoring $\hat{y}_{ui} = \mathbf{h}_u^T \mathbf{h}_i$. All analyses below are conducted on the *training split* to avoid test leakage.

To quantify *partial alignment* on social ties, we construct interaction-induced intent groups and measure directed overlap on each injected edge. Concretely, we learn item embeddings from the interaction graph only on the training split, excluding social links to avoid social leakage, and cluster items into K_c intent groups using K-means [13, 23] with multiple random initializations, selecting the partition with the lowest within-cluster SSE. Each user's intent set

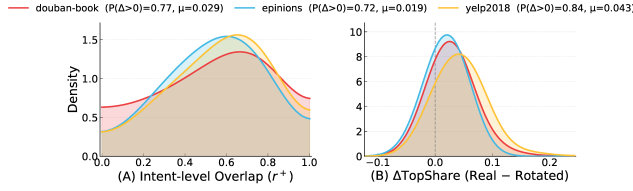


Figure 2: r^+ distribution (A) and coordinate concentration of anchor-relative deviation measured by $\Delta\text{TopShare}$ (B; Eq. (4)).

C_u is defined as the set of cluster IDs of items that user u interacted with; implementation details are provided in Appendix D.

For a directed tie $u \leftarrow v$, we define directed intent overlap as

$$r_{u \leftarrow v}^+ \triangleq \frac{|C_u \cap C_v|}{|C_v|} \in [0, 1]. \quad (1)$$

This directed form matches social injection: $r_{u \leftarrow v}^+$ measures the fraction of v 's interaction-induced intents that are also covered by u , while $1 - r_{u \leftarrow v}^+$ serves as an operational proxy for the uncovered, potentially off-anchor mass injected from v to u . A value close to 1 indicates high coverage, a value close to 0 indicates little coverage, and a mid-range value indicates partial alignment. When computing r^+ , we discard ties with $|C_v| = 0$, i.e., where v has no training interactions.

Figure 2(A) shows that r^+ concentrates in the mid-range across datasets rather than only near the two extremes. This suggests that most social ties are neither fully aligned nor fully mismatched: a single neighbor often overlaps with the target user on part of its interaction-induced intents while carrying uncovered factors at the same time. Therefore, social injection is a within-message mixture problem, where an injected message can contain both anchor-supported and off-anchor components. To ensure this phenomenon is not an artifact of clustering or degree effects, we compute intents from interactions only and repeat the analysis with multiple random seeds and a Real-vs.-Null shuffling protocol; details are reported in Appendix C and Tables 6–7.

2.2 Coordinate Concentration of Anchor-Relative Deviation

We next examine whether the off-anchor part of a social carrier is uniformly spread or concentrated on a subset of coordinates. This analysis does not assume that each embedding coordinate has an independent semantic meaning. Instead, it only tests a weaker condition: in the learned basis, whether the anchor-relative deviation has non-uniform coordinate-wise squared contribution. If so, component-wise regulation can be more expressive than scalar edge-wise scaling.

For a directed tie $u \leftarrow v$, let $\mathbf{x}_{u \leftarrow v}^{(l)} \in \mathbb{R}^D$ denote the social carrier vector injected from v to u at layer l , i.e., the per-neighbor representation of v consumed by the social aggregator. In our implementation, $\mathbf{x}_{u \leftarrow v}^{(l)}$ corresponds to the neighbor social state $\mathbf{s}_v^{(l)}$ before aggregation in Eq. (7); its relation to the aggregated decomposition follows from projection linearity, as discussed in Appendix A. We set the anchor as $\mathbf{a}_u^{(l)} \triangleq \mathbf{m}_{u,l}^{(l)}$, where $\mathbf{m}_{u,l}^{(l)}$ is an interaction-view message computed from $\mathcal{G}_l$ only. We decompose $\mathbf{x}_{u \leftarrow v}^{(l)}$ into

an anchor-parallel component and an orthogonal residual:

$$\mathbf{x}_{\parallel} = \frac{(\mathbf{x}_{u \leftarrow v}^{(l)})^\top \mathbf{a}_u^{(l)}}{\|\mathbf{a}_u^{(l)}\|_2^2 + \varepsilon} \mathbf{a}_u^{(l)}, \quad \boldsymbol{\delta} = \mathbf{x}_{u \leftarrow v}^{(l)} - \mathbf{x}_{\parallel}. \quad (2)$$

Here $\boldsymbol{\delta} \in \mathbb{R}^D$ is the anchor-relative deviation vector, and δ_k denotes the scalar value of its k -th coordinate. We call it a deviation because it is orthogonal to the interaction-grounded anchor, not because it is assumed to be pure noise.

For coordinate k , the normalized squared contribution is $\delta_k^2 / \|\boldsymbol{\delta}\|_2^2$, i.e., the fraction of the total squared deviation contributed by this coordinate. Given a budget ratio ρ , let $\mathcal{K}_\rho(\boldsymbol{\delta})$ be the indices of the top- $\lceil \rho D \rceil$ coordinates ranked by $|\delta_k|$. We define

$$\text{TopShare}(\boldsymbol{\delta}) \triangleq \frac{\sum_{k \in \mathcal{K}_\rho(\boldsymbol{\delta})} \delta_k^2}{\|\boldsymbol{\delta}\|_2^2}. \quad (3)$$

To test whether the concentration depends on the learned coordinate system, we compare the original deviation with a Haar-randomly rotated version. Let $\mathbf{Q} \in \mathbb{R}^{D \times D}$ be a random orthogonal matrix and let $\boldsymbol{\delta}_{\text{rot}} = \mathbf{Q}\boldsymbol{\delta}$. Since $\mathbf{Q}$ preserves the ℓ_2 norm but randomizes coordinate axes, concentration tied to the learned basis should decrease after rotation. We compute

$$\Delta\text{TopShare} \triangleq \text{TopShare}(\boldsymbol{\delta}) - \text{TopShare}(\boldsymbol{\delta}_{\text{rot}}). \quad (4)$$

Here $\mathbf{Q}$ is sampled as a Haar-random orthogonal matrix, preserving the ℓ_2 norm while randomizing coordinate axes; details are provided in Appendix E. A positive $\Delta\text{TopShare}$ means that the original deviation is more coordinate-concentrated than its randomly rotated counterpart under the same coordinate budget.

Figure 2(B) shows a consistent positive shift across datasets. At $\rho = 0.3$, $P(\Delta\text{TopShare} > 0)$ is 0.72, 0.84, and 0.77, with mean shifts 0.019, 0.043, and 0.029 on Epinions, Yelp2018, and Douban-Book, respectively. Thus, anchor-relative deviations are coordinate-concentrated in the learned basis. For example, one Douban-Book tie yields $\text{TopShare}(\boldsymbol{\delta}) = 0.905$ versus $\text{TopShare}(\mathbf{Q}\boldsymbol{\delta}) = 0.779$ using 20 of 64 coordinates. This trend remains stable for $\rho = 0.1$ –0.5.

Implication: why scalar edge reweighting is insufficient.

The two analyses above imply a mismatch between edge-space control and within-message mixture. Under partial alignment, a social carrier can be written as $\mathbf{x} = \mathbf{x}_{\parallel} + \boldsymbol{\delta}$, where $\mathbf{x}_{\parallel}$ is anchor-aligned and $\boldsymbol{\delta}$ is anchor-relative deviation. A scalar edge weight α gives $\alpha\mathbf{x} = \alpha\mathbf{x}_{\parallel} + \alpha\boldsymbol{\delta}$, which cannot attenuate the off-anchor deviation without also shrinking the aligned component. Moreover, because $\boldsymbol{\delta}$ is coordinate-concentrated in the learned basis, component-wise regulation can selectively control high-contribution deviation coordinates. This does not require each coordinate to be semantically disentangled; it only requires the deviation to be non-uniformly distributed across coordinates. Therefore, ODSR adopts a dimension-wise gate to softly regulate the deviation rather than deleting it or uniformly scaling the whole social message.

3 Methodology

Overview of ODSR. For each user u and layer l , ODSR uses the interaction-view message $\mathbf{m}_{u,l}^{(l)}$ as an interaction-grounded anchor, decomposes the aggregated social message into an anchor-aligned

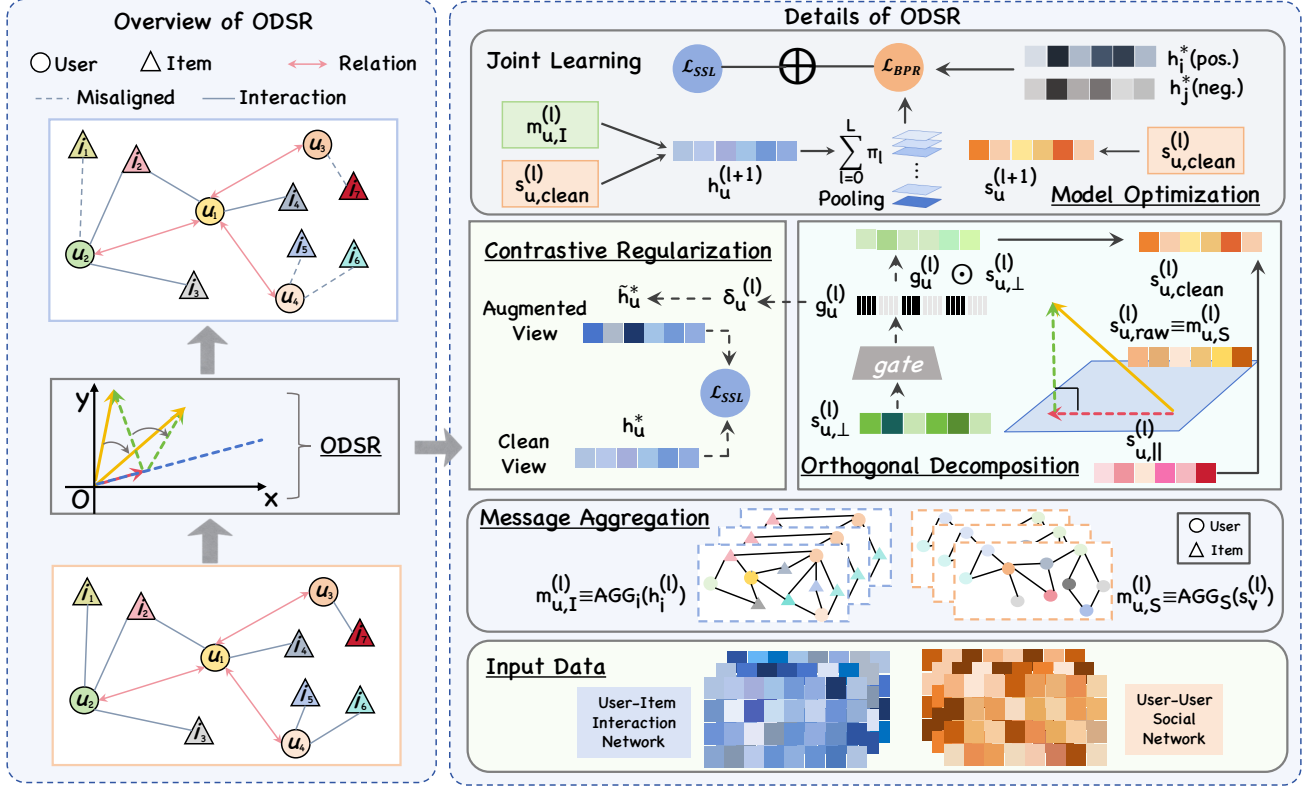


Figure 3: Architecture of ODSR. Left: Overview of ODSR, illustrating the regulation of social injection for a target user. Right: Details (read bottom-up): ODSR takes the interaction and social networks as inputs, performs dual-view aggregation, decomposes the social message into an anchor-parallel component $s_{u,||}^{(l)}$ and an orthogonal deviation $s_{u,\perp}^{(l)}$, learns a dimension-wise gate $g_u^{(l)}$ to form $s_{u,clean}^{(l)}$, and jointly optimizes ranking with deviation-aware contrastive regularization.

component and an orthogonal deviation, and learns a supervision-guided dimension-wise gate to admit only beneficial deviation coordinates for ranking. This yields Eq. (5), where the gate modulates $s_{u,\perp}^{(l)}$ rather than uniformly scaling neighbors, preserving aligned signals while regulating risky residuals. Each layer performs interaction/social aggregation, anchor-orthogonal decomposition, and deviation gating, with a deviation-aware contrastive regularizer stabilizing the control boundary. Formally, ODSR updates the user embedding at layer l as:

$$h_u^{(l+1)} = m_{u,I}^{(l)} + \underbrace{s_{u,||}^{(l)} + g_u^{(l)} \odot s_{u,\perp}^{(l)}}_{\text{filtered social injection with control in embedding space}}, \quad (5)$$

where $m_{u,I}^{(l)}$ is the interaction-view message, computed by Eq. (6), that defines the anchor direction, $s_{u,||}^{(l)}$ and $s_{u,\perp}^{(l)}$ are the anchor-aligned component and the orthogonal deviation of the social message, $g_u^{(l)} \in (0, 1)^D$ is a dimension-wise gating vector, and $\odot$ denotes the element-wise Hadamard product.

3.1 Anchor-Orthogonal Social Decomposition

3.1.1 Dual-View Message Passing. Since social aggregation may mix task-aligned factors with off-anchor deviations, we separate an anchor message from a social carrier message so that deviation control targets the injected social component.

Given neighbors $\mathcal{N}_u^I = \{i \in \mathcal{V} \mid (u, i) \in \mathcal{E}_I\}$ and $\mathcal{N}_u^S = \{v \in \mathcal{U} \mid (u, v) \in \mathcal{E}_S\}$, we define a ranking representation $h_u^{(l)}$ and a social state $s_u^{(l)}$. These two views are coupled by injecting a *filtered* social message into the interaction-grounded update (Eq. (5)). The auxiliary state $s_u^{(l)}$ acts as a stable carrier for social propagation, decoupling social mixing from the ranking embedding update so that regulation operates explicitly in the embedding space.

Interaction View. We follow standard LightGCN propagation [8] on the interaction graph:

$$m_{u,I}^{(l)} = \sum_{i \in \mathcal{N}_u^I} \frac{1}{\sqrt{d_u^l d_i^l}} h_i^{(l)}. \quad (6)$$

where d_u^l and d_i^l denote the degrees of user u and item i in the interaction graph. This message is naturally aligned with the ranking objective and defines the task-grounded anchor direction.

Social View. We maintain an auxiliary social state $\mathbf{s}_u^{(l)}$ initialized by $\mathbf{s}_u^{(0)} = \mathbf{h}_u^{(0)}$ [18]. The raw social message is computed by row-normalized aggregation:

$$\mathbf{m}_{u,S}^{(l)} = \sum_{v \in N_u^S} \frac{1}{d_u^S} \mathbf{s}_v^{(l)}. \quad (7)$$

where d_u^S denotes the number of social neighbors aggregated for u , i.e., the corresponding degree in the directed social graph.

3.1.2 Anchor-Orthogonal Decomposition. We decompose the aggregated social message relative to an anchor direction grounded in interactions. As illustrated in Figure 3 (right), this step explicitly separates the socially propagated carrier into an anchor-aligned part and a deviation part, so that subsequent gating targets the risky component.

$$\mathbf{a}_u^{(l)} \triangleq \mathbf{m}_{u,I}^{(l)}. \quad (8)$$

We project $\mathbf{m}_{u,S}^{(l)}$ onto $\text{span}(\mathbf{a}_u^{(l)})$ and its orthogonal complement:

$$\mathbf{s}_{u,\parallel}^{(l)} = \frac{(\mathbf{m}_{u,S}^{(l)})^\top \mathbf{a}_u^{(l)}}{\|\mathbf{a}_u^{(l)}\|_2^2 + \epsilon} \mathbf{a}_u^{(l)}, \quad \mathbf{s}_{u,\perp}^{(l)} = \mathbf{m}_{u,S}^{(l)} - \mathbf{s}_{u,\parallel}^{(l)}. \quad (9)$$

3.2 Dimension-wise Gating in Embedding Space

3.2.1 Deviation Admission with a Gating Vector. Existing edge-wise scaling uniformly affects both aligned and off-anchor components and thus cannot resolve their mixture. To address this, ODSR regulates $\mathbf{s}_{u,\perp}^{(l)}$ via a dimension-wise gating vector $\mathbf{g}_u^{(l)}$. Motivated by our analysis that off-anchor deviation concentrates on a subset of dimensions in the learned embedding basis, we construct a compact gate context from the anchor and deviation vectors, their feature-wise compatibility, and two lightweight geometric evidence scalars $e_{\cos}^{(l)}$ and $e_{\text{norm}}^{(l)}$ appended as scalar features:

$$e_{\cos}^{(l)} = \frac{(\mathbf{m}_{u,S}^{(l)})^\top \mathbf{a}_u^{(l)}}{\|\mathbf{m}_{u,S}^{(l)}\|_2 \|\mathbf{a}_u^{(l)}\|_2 + \epsilon}, \quad e_{\text{norm}}^{(l)} = \tanh(\|\mathbf{m}_{u,S}^{(l)}\|_2). \quad (10)$$

3.2.2 Gate Mechanism: Lightweight Gating Vector. Let the gate input be:

$$\mathbf{v}_u^{(l)} = \left[\mathbf{a}_u^{(l)}; \mathbf{s}_{u,\perp}^{(l)}; \mathbf{a}_u^{(l)} \odot \mathbf{s}_{u,\perp}^{(l)}; e_{\cos}^{(l)}; e_{\text{norm}}^{(l)} \right]. \quad (11)$$

where the Hadamard interaction captures feature-wise compatibility between the anchor direction and the deviation. The gate is parameterized by a *single* affine transformation with a sigmoid:

$$\mathbf{g}_u^{(l)} = \sigma(\mathbf{W}_g \mathbf{v}_u^{(l)} + \mathbf{b}_g). \quad (12)$$

Here $\mathbf{W}_g$ maps the concatenated context to $\mathbb{R}^D$. The gate parameters are optimized jointly with the ranking objective, so deviation admission is directly shaped by ranking signals. This lightweight design keeps the admission rule expressive at the dimension level while reducing overfitting risk under weak homophily.

3.2.3 Filtered Social Injection and Social State Update. The filtered social injection is:

$$\mathbf{s}_{u,\text{clean}}^{(l)} = \mathbf{s}_{u,\parallel}^{(l)} + \mathbf{g}_u^{(l)} \odot \mathbf{s}_{u,\perp}^{(l)}. \quad (13)$$

We update the auxiliary social state by setting

$$\mathbf{s}_u^{(l+1)} = \mathbf{s}_{u,\text{clean}}^{(l)}. \quad (14)$$

This update serves as a stable carrier for social propagation, while ranking follows Eq. (5).

3.3 Risk-Aware Contrastive Regularization

We further stabilize training by contrasting a clean view with a *propagation-augmented* view constructed via deviation-guided perturbation. Rather than using generic isotropic noise, we align the augmentation with the model's suppression pattern to regularize the gate-induced control boundary, making the learned admission less sensitive to small but adversarial off-anchor shifts.

Layer-wise Deviation Perturbation. We denote the stop-gradient operator as $\text{sg}(\cdot)$, which blocks gradients from flowing through its argument. After computing $\mathbf{g}_u^{(l)}$ and $\mathbf{s}_{u,\perp}^{(l)}$, we normalize the deviation direction and apply an element-wise magnitude mask (scaled by ϵ) on dimensions that the gate tends to suppress:

$$\delta_u^{(l)} = \left(\epsilon \cdot (1 - \text{sg}(\mathbf{g}_u^{(l)})) \right) \odot \frac{\mathbf{s}_{u,\perp}^{(l)}}{\|\mathbf{s}_{u,\perp}^{(l)}\|_2 + \epsilon}. \quad (15)$$

where ϵ controls the perturbation magnitude. We add it to the user update to construct the augmented view:

$$\tilde{\mathbf{h}}_u^{(l+1)} \triangleq \mathbf{h}_u^{(l+1)} + \delta_u^{(l)}. \quad (16)$$

We stop gradients through $\mathbf{g}_u^{(l)}$ in Eq. (15) to avoid objective conflict between gating and contrastive regularization.

InfoNCE Objective. Let $\mathbf{h}_u^*$ and $\tilde{\mathbf{h}}_u^*$ denote the pooled user representations of the clean and augmented views, computed using the layer pooling in Eq. (18). We minimize the InfoNCE loss [7, 33, 40]:

$$\mathcal{L}_{SSL} = - \sum_{u \in \mathcal{B}} \log \frac{\exp(\text{sim}(\mathbf{h}_u^*, \tilde{\mathbf{h}}_u^*)/\tau)}{\sum_{v \in \mathcal{B}} \exp(\text{sim}(\mathbf{h}_u^*, \tilde{\mathbf{h}}_v^*)/\tau)}, \quad (17)$$

where $\mathcal{B}$ is the sampled mini-batch of users, $\text{sim}(\cdot, \cdot)$ is cosine similarity, and τ is the temperature.

3.4 Model Optimization and Complexity

Learnable Layer Pooling. Following common practice in graph recommendation, we aggregate embeddings across layers. ODSR adopts a learnable weighted pooling:

$$\mathbf{h}_u = \sum_{l=0}^L \pi_l \mathbf{h}_u^{(l)}, \quad \mathbf{h}_i^* = \sum_{l=0}^L \pi_l \mathbf{h}_i^{(l)}, \quad \boldsymbol{\pi} = \text{softmax}(\boldsymbol{\alpha}), \quad (18)$$

where $\boldsymbol{\alpha}$ are learnable layer scores and $\boldsymbol{\pi}$ are normalized weights. This allows the model to adaptively emphasize layers with more informative signals under social injection.

Optimization Objective. We optimize BPR loss [28] on sampled triples (u, i, j) :

$$\mathcal{L}_{BPR} = \sum_{(u,i,j) \in \mathcal{D}_{train}} -\ln \sigma(\mathbf{h}_u^{*\top} \mathbf{h}_i^* - \mathbf{h}_u^{*\top} \mathbf{h}_j^*), \quad (19)$$

and the full objective is:

$$\mathcal{L} = \mathcal{L}_{BPR} + \lambda_{ssl} \mathcal{L}_{SSL} + \lambda_{reg} \|\Theta\|_2^2. \quad (20)$$

Training Procedure. At each iteration, we sample a mini-batch of triples, perform full-graph sparse propagation with anchor computation, orthogonal decomposition, dimension-wise gating, and deviation-guided augmentation, and then optimize the joint objective in Eq. (20). The model is updated by back-propagation with the BPR loss and the deviation-aware contrastive loss, following the standard training paradigm of graph recommendation models.

Complexity Analysis. Sparse message passing costs $O(L(|\mathcal{E}_I| + |\mathcal{E}_S|)D)$. The gate applies one dense affine map per user per layer, costing $O(L|\mathcal{U}|D D_{\text{in}})$ with $D_{\text{in}} = 3D + 2$ from Eq. (11), i.e., $D_{\text{in}} = O(D)$. In practice, sparse propagation over $|\mathcal{E}_I| + |\mathcal{E}_S|$ typically dominates and the gate adds a moderate overhead; in typical social recommendation graphs where $|\mathcal{E}_I| + |\mathcal{E}_S| \gg |\mathcal{U}|$, sparse propagation remains the main cost. Memory usage is $O((|\mathcal{U}| + |\mathcal{V}|)D)$ for storing embeddings. Consequently, for a fixed embedding size, the complexity remains linear in the graph size in our setting.

4 Experiments and Analysis

4.1 Experimental Settings

4.1.1 Datasets. We conduct experiments on three social recommendation benchmarks: Epinions, Yelp2018, and Douban-Book. For Epinions, we adopt the processed versions from recent studies [14, 46] to match state-of-the-art benchmark settings. For Yelp2018, we adopt the processed version from Yang et al. [45]. For Douban-Book, we use the version provided by SelfRec [51]. We follow the evaluation protocol and data splits provided by the corresponding processed benchmarks. Statistics are reported in Table 1.

We compare ODSR with 15 strong baselines spanning three representative classes of social recommendation (SR) methods widely adopted in recent literature: (1) *General graph collaborative filtering*: LightGCN (SIGIR'20)[8], LightGCN-S (as in GBSR)[47], SimGCL (SIGIR'22)[50], and DCCF (SIGIR'23)[27]; (2) *Self-supervised and robust SR*: SEPT (KDD'21)[48], GDMSR (WWW'23)[26], MADM (WSDM'24)[22], SHaRe (WWW'24)[12], GBSR (KDD'24)[47], SGIL (SIGIR'25)[46], DSL (IJCAI'23)[35], and SSD-ICGA (KDD'24)[30]; (3) *Generative refiners*: RecDiff (CIKM'24)[17], ARD-SR (WWW'25)[31], and RecFlow (ICML'25)[9]. We omit classic MF-based and early SR baselines, such as NGCF[36], TrustMF[44], GraphRec[3], and DiffNet++[41], as recent diffusion-based SR refiners[17, 31] consistently outperform them on standard benchmarks.

4.1.2 Evaluation Metrics. We evaluate ranking quality using Recall@K and NDCG@K [11] with $K \in \{10, 20\}$, abbreviated as R@K and N@K in the tables. We adopt the full-ranking protocol [47] by ranking all items per user except training interactions. We report the mean results over 5 independent runs and conduct paired t -tests with $p < 0.05$ for statistical significance.

Implementation Details. All models use 64-dimensional embeddings with Xavier initialization, Adam with a learning rate of 0.001 and a batch size of 2048, three graph propagation layers, BPR with one uniformly sampled negative item per positive interaction, and early stopping after 50 non-improving validation epochs. For ODSR, we set $\tau = 0.2$ and tune $\lambda_{\text{reg}} \in \{10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}\}$ and $\lambda_{\text{ssl}}, \epsilon \in [0.1, 0.8]$ with a step size of 0.1.

Table 1: Dataset statistics.

Dataset	# Users	# Items	# Interactions	# Social links
Epinions	18,202	47,449	298,173	381,559
Yelp2018	45,919	45,538	1,183,610	709,459
Douban-Book	13,025	22,348	598,420	169,150

4.2 Results and Analysis

4.2.1 Overall Performance Comparison. The performance of all methods on three datasets is reported in Table 2, where * indicates statistically significant improvement over the strongest baseline. We highlight four findings. First, incorporating social information is not uniformly beneficial: naive social propagation, such as LightGCN-S, yields only marginal gains over LightGCN and can be inferior to stronger robust/SSL baselines, suggesting that directly injecting social signals may introduce mismatch under weak preference homophily. Second, robust social recommendation methods, including GBSR, SGIL, MADM, and SSD-ICGA, are consistently among the strongest competitors, which supports the view that social noise and partial alignment are pervasive in open platforms and must be handled beyond vanilla aggregation. Third, generative refiners, such as RecDiff, ARD-SR, and RecFlow, are competitive but do not dominate under the same protocol, indicating that multi-step refinement is not strictly necessary to obtain strong robustness if the injected message is properly regulated. Finally, ODSR achieves the best results on all datasets, improving over the strongest baseline by +6.26% NDCG@20 on Epinions, +2.29% on Yelp2018, and +2.71% on Douban-Book; the consistent gains across datasets validate that selectively controlling off-anchor deviations within the aggregated social message complements edge-space denoising and improves ranking robustness.

Mechanism: deviation ratio reduction. We verify component-wise deviation regulation by decomposing each layer-0 social carrier into $m_{\parallel}$ and $m_{\perp}$ and comparing the deviation ratio $\|m_{\perp}\|_2 / \|m\|_2$ before and after gating. Across all datasets, the post-gating ratio is consistently smaller than the pre-gating ratio, indicating that ODSR reduces off-anchor deviation without collapsing the aligned component. As complementary evidence, Appendix B reports an inference-time all-pass intervention showing that removing deviation control hurts NDCG@20.

4.2.2 Ablation Study. We conduct ablation studies to validate the contributions of key components in ODSR. Results are reported in Table 3. All variants use the same backbone and training protocol as ODSR, and we only modify the specified component. We report both Recall@20 and NDCG@20 on the three datasets.

Effect of anchor-orthogonal decomposition and gating. Removing the decomposition module is one of the most damaging ablations, confirming the importance of explicitly separating aligned and orthogonal components. This indicates that directly injecting the aggregated social message is insufficient under partial alignment, because aligned and off-anchor components are mixed inside the same message. The all-pass gate variant also performs worse than ODSR, showing that simply preserving the full orthogonal

Table 2: Overall performance comparison using Recall and NDCG at 10 and 20. Best is bold, second best is underlined. * indicates statistically significant improvement over the strongest baseline with $p < 0.05$.

Model	Epinions				Yelp2018				Douban-Book			
	R@10	N@10	R@20	N@20	R@10	N@10	R@20	N@20	R@10	N@10	R@20	N@20
LightGCN	0.0432	0.0314	0.0675	0.0385	0.0416	0.0351	0.0689	0.0449	0.0983	0.1154	0.1478	0.1242
LightGCN-S	0.0452	0.0323	0.0696	0.0394	0.0431	0.0364	0.0709	0.0464	0.1010	0.1177	0.1500	0.1265
SEPT	0.0457	0.0341	0.0724	0.0416	0.0412	0.0347	0.0696	0.0448	0.1027	0.1204	0.1547	0.1314
DCCF	0.0425	0.0310	0.0665	0.0380	0.0361	0.0303	0.0690	0.0448	0.0985	0.0998	0.1428	0.1191
SimGCL	0.0479	0.0356	0.0740	0.0435	0.0431	0.0361	0.0711	0.0457	0.1125	0.1255	0.1615	0.1360
GDMSR	0.0472	0.0350	0.0732	0.0426	0.0420	0.0356	0.0694	0.0453	0.0990	0.1165	0.1483	0.1248
DSL	0.0478	0.0355	0.0738	0.0431	0.0424	0.0359	0.0698	0.0455	0.0995	0.1172	0.1486	0.1251
SHaRe	0.0490	0.0354	0.0771	0.0438	0.0416	0.0373	0.0732	0.0478	0.1005	0.1182	0.1491	0.1255
SSD-ICGA	0.0485	0.0360	0.0752	0.0440	0.0430	0.0365	0.0712	0.0465	0.0998	0.1178	0.1488	0.1252
GBSR	0.0515	0.0372	0.0795	0.0458	0.0462	0.0395	0.0768	0.0498	0.1030	0.1215	0.1518	0.1285
MADM	0.0535	0.0386	0.0835	0.0472	0.0445	0.0367	0.0755	0.0468	<u>0.1313</u>	0.1505	0.1752	0.1563
SGIL	<u>0.0542</u>	<u>0.0392</u>	<u>0.0843</u>	<u>0.0479</u>	<u>0.0484</u>	<u>0.0412</u>	<u>0.0797</u>	<u>0.0524</u>	0.1300	<u>0.1563</u>	<u>0.1808</u>	<u>0.1625</u>
RecDiff	0.0528	0.0380	0.0820	0.0465	0.0438	0.0376	0.0735	0.0481	0.1008	0.1185	0.1495	0.1258
ARD-SR	0.0520	0.0375	0.0810	0.0460	0.0445	0.0380	0.0746	0.0486	0.1022	0.1205	0.1508	0.1269
RecFlow	0.0522	0.0378	0.0815	0.0462	0.0452	0.0388	0.0758	0.0495	0.1025	0.1208	0.1511	0.1272
ODSR	0.0567*	0.0416*	0.0885*	0.0509*	0.0494*	0.0420*	0.0817*	0.0536*	0.1325*	0.1596*	0.1867*	0.1669*

Table 3: Ablation study using Recall@20 and NDCG@20.

Variant	Epinions		Yelp2018		Douban-Book	
	R@20	N@20	R@20	N@20	R@20	N@20
ODSR	0.0885	0.0509	0.0817	0.0536	0.1867	0.1669
w/o Decomposition	0.0758	0.0438	0.0668	0.0435	0.1565	0.1395
w/ All-pass Gate	0.0785	0.0454	0.0742	0.0484	0.1715	0.1528
w/ Scalar Gate	0.0805	0.0466	0.0768	0.0501	0.1760	0.1570
w/o $e_{\cos}$	0.0817	0.0470	0.0802	0.0528	0.1835	0.1649
w/o e_{norm}	0.0822	0.0473	0.0799	0.0527	0.1835	0.1644
w/o both cues	0.0818	0.0475	0.0799	0.0527	0.1831	0.1637
w/ Isotropic SSL	0.0762	0.0441	0.0710	0.0463	0.1680	0.1498
w/o SSL	0.0695	0.0402	0.0640	0.0415	0.1620	0.1445

deviation introduces noisy social factors. Replacing the dimension-wise gate with a scalar gate improves over all-pass injection but remains inferior to ODSR, which confirms that a single global coefficient cannot selectively regulate different deviation coordinates.

Effect of geometric evidence cues. We further examine the two lightweight evidence cues used in the gate input. Removing either $e_{\cos}$ or e_{norm} causes consistent degradation compared with ODSR. The variant without both cues also underperforms the full model, suggesting that the cosine cue and norm cue provide complementary information for estimating whether the social deviation should

be admitted. Although these variants are stronger than removing decomposition or SSL entirely, their gaps to ODSR show that geometric evidence helps the gate make more reliable dimension-wise admission decisions.

Effect of deviation-aware contrastive regularization. Replacing deviation-guided SSL with isotropic SSL leads to clear performance drops, especially on Epinions and Douban-Book. This suggests that generic random perturbations are less suitable for modeling the specific risk brought by off-anchor social deviations. Removing SSL entirely causes the worst or near-worst performance among the SSL-related variants, verifying that the contrastive objective stabilizes the learned deviation-control boundary and improves robustness against imperfect social signals.

4.2.3 Performance under Different Sparsity Levels. To evaluate performance under data scarcity, we rank users in ascending order by their training interaction degrees and divide them into four groups: < 10 , $10-19$, $20-50$, and > 50 . We compare ODSR with three representative and strong competitors, SGIL, MADM, and LightGCN-S, and report NDCG@10 in Figure 4. Two trends are consistent across datasets. First, all methods improve as users become less sparse, which is consistent with the intuition that denser supervision yields more reliable preference estimates. Second, ODSR shows clear advantages in sparse regimes, where social signals are both valuable and risk-prone. Notably, the gain is most pronounced for the cold-start group (degree < 10), where interaction-only supervision is weakest and unregulated social injection is more likely to

Table 4: Empirical efficiency comparison with SGIL under the same RTX 4090 environment.

Dataset	Time / Epoch (s)		Peak Mem. (MB)	
	SGIL	ODSR	SGIL	ODSR
Epinions	36.79	4.19	2628	846
Yelp2018	435.92	24.26	5131	1514
Douban-Book	40.50	7.94	2440	504

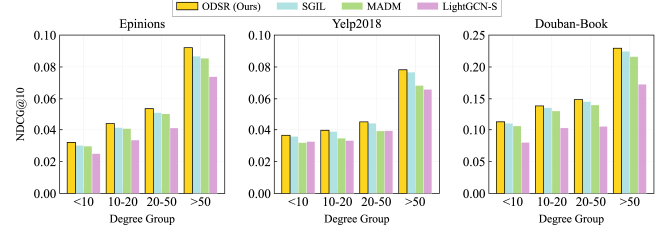
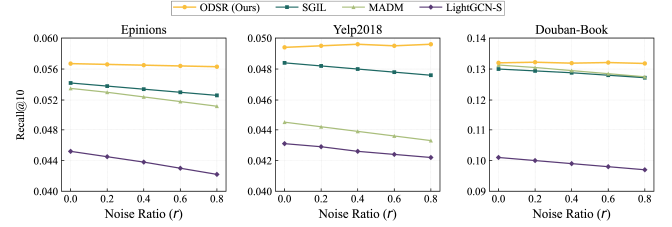
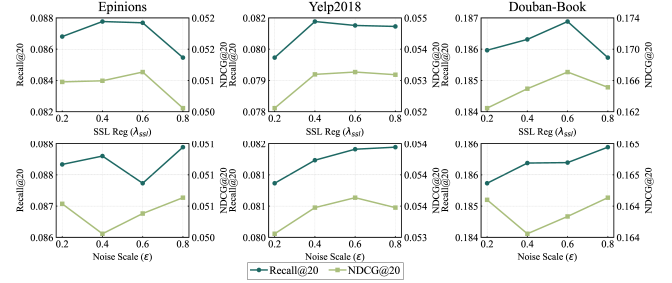
introduce off-anchor deviations. We select SGIL [46], MADM [22], and LightGCN-S [47] as representative baselines covering invariant learning, adaptive edge pruning, and naive social propagation.

4.2.4 Robustness to Social Noise. To substantiate robustness to social noise, we contaminate the social network by injecting different proportions of fake social connections while keeping the user-item interaction graph unchanged. Such fake ties mainly introduce off-anchor deviations into the aggregated social message, which ODSR can suppress at the component level.

Given a noise ratio r , we add $r \cdot |\mathcal{E}_S|$ fake social links into the observed social graph, where $|\mathcal{E}_S|$ denotes the number of original social links. Fake links are uniformly sampled from all user pairs excluding self-loops and existing edges, and duplicates are removed to avoid multi-edges. For directed social graphs, we sample ordered user pairs. For undirected graphs, we add reciprocal links for each sampled pair. For each r , we corrupt the social graph before training and retrain all compared methods from scratch under the same settings. We compare ODSR with SGIL, MADM, and LightGCN-S, and report Recall@10 under $r \in \{0.0, 0.2, 0.4, 0.6, 0.8\}$ in Figure 5. We observe three consistent patterns. (1) Performance drops as noise increases, indicating sensitivity to spurious social links. (2) ODSR achieves the highest accuracy and degrades less with larger r , showing stronger robustness to fake ties. (3) At low-to-moderate noise, degradation is mild because informative edges remain and normalization dilutes injected links. Overall, ODSR is the most stable under noise, likely due to within-embedding deviation regulation.

4.2.5 Efficiency Analysis. Table 4 reports the measured training time and peak memory of ODSR and SGIL under the same RTX 4090 environment. ODSR is consistently faster and uses less GPU memory on all three datasets, suggesting that the additional embedding-space regulation is lightweight in practice.

4.2.6 Hyper-parameter Sensitivity Analysis. We investigate the sensitivity of ODSR to two key hyperparameters: the self-supervised regularization weight λ_{ssl} and the augmentation noise scale ϵ . For each configuration, we repeat the experiment 5 times with different random seeds and report averaged Recall@20 and NDCG@20. Unless otherwise specified, all remaining hyperparameters are fixed to their dataset-specific best values obtained by grid search, and we vary one parameter in $[0.0, 1.0]$ while fixing the other at its best value. As illustrated in Figure 6, performance improves with increasing λ_{ssl} and peaks around 0.2 to 0.6 depending on the dataset, while an overly large weight above 0.8 tends to overshadow the BPR loss. Regarding ϵ , ODSR is stable in 0.2 to

**Figure 4: Performance (NDCG@10) across sparsity levels.****Figure 5: Performance (Recall@10) across noise ratios.****Figure 6: Sensitivity to hyperparameters λ_{ssl} and ϵ .**

0.8 with optimal performance clustering around 0.4 to 0.6. The larger optimal ϵ in ODSR suggests that perturbations guided by orthogonal deviations are structurally informative, allowing stronger robustness training without collapsing representations.

4.2.7 Case Study. To provide an understanding of selective denoising, we perform a case study on two randomly sampled users, one from Epinions and one from Douban-Book, and visualize their layer-0 social injections in Figure 7. Blue edges indicate the anchor-aligned component, and yellow edges indicate the anchor-orthogonal deviation. For deviation edges, we annotate the magnitude before and after gating using w_{pre} and w_{post} .

Aligned components are not annotated since ODSR mainly regulates the deviation part through the gating mechanism. Across both examples, the gate consistently preserves anchor-aligned signals while selectively attenuating anchor-orthogonal deviations. Specifically, a reduction from w_{pre} to w_{post} indicates that the corresponding off-anchor deviation is only partially admitted after gating. This highlights a direction-aware effect: the gate regulates within-embedding mixtures by modulating only the deviation component, without uniformly reweighting neighbors. The annotated

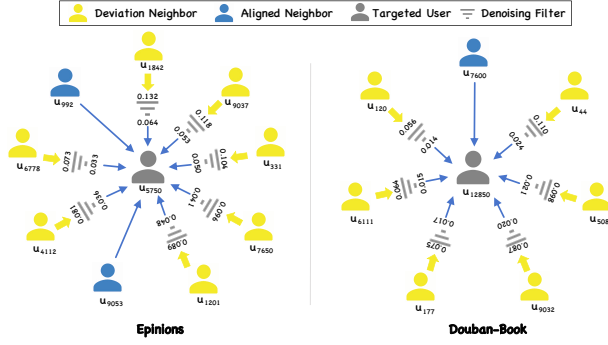


Figure 7: Case study for the deviation-gating effect of ODSR.

$w_{pre} \rightarrow w_{post}$ on each deviation edge illustrates the attenuation of the orthogonal deviation after gating.

5 Related Work

5.1 GNN-based Social Recommendation and Self-Supervision

Early social recommendation primarily incorporated trust relations and social ties into matrix factorization models [6, 10, 20, 21]. Recent methods adopt GNN-style message passing to propagate high-order collaborative and social signals, where simplified propagation (e.g., LightGCN [8] and SocialLGN [18]) relies on neighborhood aggregation and representation smoothing. Such smoothing can be effective when neighborhood signals are preference-consistent, but may amplify mismatched information under weak preference homophily, which is common in open social platforms [41]. Self-supervised learning further improves generalization by constructing multiple views and enforcing cross-view consistency, e.g., via stochastic perturbations in structure or embedding space (SGL [40], SimGCL [50], LightGCL [1]). Neural influence diffusion is another representative line that explicitly models how social signals propagate to shape preference [42]. However, existing SSL designs and aggregation mechanisms provide limited control over *how* the injected social representation deviates relative to task-grounded preference directions, motivating direction-aware regulation under social aggregation in our embedding-space setting [49, 51].

5.2 Edge-Space Robustification: Structure Refinement and Scalar Reweighting

Robust social recommendation has been widely studied under noisy or weakly-homophilous social graphs. A dominant line of work intervenes in the *edge space* by learning or modifying the social graph structure, such as rewiring in SHaRe [12] and counterfactual refinement strategies [14], often coupled with denoising objectives [30, 35]. Another line seeks robust or invariant signals to reduce sensitivity to mismatch, e.g., invariant learning in SGIL [46]. Many approaches also assign scalar confidence/attention scores to filter edges or reweight neighbor messages (e.g., GDMSR [26], GBSR [47], MADM [22]). While effective, scalar reweighting is direction-invariant: it uniformly scales an aggregated message and cannot selectively regulate different components when task-aligned

signals and off-anchor deviations coexist within the same vector. This limitation motivates a complementary intervention that operates *within* the message representation to enable component-wise, controllable information exchange.

5.3 Generative Refinement and Disentangled Recommendation

Generative approaches have recently shown promise in refining noisy graph structures. For instance, diffusion-based methods like RecDiff [17] and Graph Diffusive Learning [15] employ forward-reverse processes to generate reliable representations, while ARD-SR [31] focuses on refining the social adjacency matrix. Similarly, flow-matching techniques have been introduced to model continuous dynamics for efficient denoising [9]. However, these methods typically rely on iterative generation steps, introducing additional computational complexity. Parallel to this, disentangled learning aims to separate latent factors behind user behaviors [27]. Pioneering works like DGCF [38] disentangle user intents on interaction graphs. In the social domain, recent studies have extended this idea to distinguish between social influence and personal interest [16, 52], or to separate multi-faceted social relations [2, 29, 43]. While some works explore orthogonality constraints to ensure factor independence [39], our ODSR distinctively focuses on the *within-embedding* mixture induced by partial alignment, achieving anchor-relative regulation via a lightweight, single-step orthogonal gate without complex generative procedures.

6 Conclusion

We study robust social recommendation under partial alignment and show that social aggregation often induces a mixture within a message, where task-aligned signals and off-anchor deviations coexist and cannot be reliably addressed by edge-wise scalar reweighting. To address this issue, we propose ODSR, an embedding-space framework that performs Orthogonal Decomposition with an anchor and selectively admits the orthogonal deviation via a dimension-wise gate learned from ranking signals. We further introduce a deviation-aware contrastive regularizer that perturbs representations along orthogonal directions to stabilize the learned control boundary. Experiments on three datasets demonstrate that ODSR consistently outperforms strong baselines and improves robustness under sparse interactions and increasing social noise. Moreover, ODSR offers an explicit and interpretable mechanism to quantify and regulate socially injected deviations in embedding space, complementing existing edge-space denoising strategies. Overall, these results verify that controlling within-message deviations is a key robustness lever for social recommendation, yielding consistent gains without modifying the underlying social graph.

Acknowledgments

The authors acknowledge the financial support from the National Natural Science Foundation of China (Grant Nos. 62276171, 62476173, 62532007), the Shenzhen Fundamental Research-General Project (Grant Nos. JCYJ20240813141503005, JCYJ20240813142610014, ZDCY-20250901110940006), and the Major Special Project for Philosophy and Social Sciences Research of the Ministry of Education (Grant No. 2025JZDZ010). Hao Liao is the corresponding author.

References

- [1] Xuheng Cai, Chao Huang, Lianghao Xia, and Xubin Ren. 2023. LightGCL: Simple Yet Effective Graph Contrastive Learning for Recommendation. In *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*. OpenReview.net. <https://openreview.net/forum?id=FKXVK9dyMM>
- [2] Chong Chen, Min Zhang, Yiqun Liu, and Shaoping Ma. 2019. Social Attentional Memory Network: Modeling Aspect- and Friend-Level Differences in Recommendation. In *Proceedings of the Twelfth ACM International Conference on Web Search and Data Mining* (Melbourne VIC, Australia) (*WSDM '19*). Association for Computing Machinery, New York, NY, USA, 177–185. doi:10.1145/3289600.3290982
- [3] Wenqi Fan, Yao Ma, Qing Li, Yuan He, Yihong Eric Zhao, Jiliang Tang, and Dawei Yin. 2019. Graph Neural Networks for Social Recommendation. In *The World Wide Web Conference, WWW 2019, San Francisco, CA, USA, May 13-17, 2019*, Ling Liu, Ryen W. White, Amin Mantrach, Fabrizio Silvestri, Julian J. McAuley, Ricardo Baeza-Yates, and Leila Zia (Eds.). ACM, 417–426. doi:10.1145/3308558.3313488
- [4] Mark S. Granovetter. 1973. The Strength of Weak Ties. *Amer. J. Sociology* 78, 6 (1973), 1360–1380.
- [5] Mark S. Granovetter. 1983. The Strength of Weak Ties: A Network Theory Revisited. *Sociological Theory* 1, 6 (1983), 201–233.
- [6] Guibing Guo, Jie Zhang, and Neil Yorke-Smith. 2015. TrustSVD: Collaborative Filtering with Both the Explicit and Implicit Influence of User Trust and of Item Ratings. In *Proceedings of the Twenty-Ninth AAAI Conference on Artificial Intelligence, January 25-30, 2015, Austin, Texas, USA*, Blai Bonet and Sven Koenig (Eds.). AAAI Press, 123–129. doi:10.1609/AAAI.V29I1.9153
- [7] Michael Gutmann and Apo Hyvärinen. 2010. Noise-contrastive estimation: A new estimation principle for unnormalized statistical models. In *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics, AISTATS 2010, Chia Laguna Resort, Sardinia, Italy, May 13-15, 2010 (JMLR Proceedings, Vol. 9)*, Yee Whye Teh and D. Mike Titterton (Eds.). JMLR.org, Chia Laguna Resort, Sardinia, Italy, 297–304. <http://proceedings.mlr.press/v9/gutmann10a.html>
- [8] 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, SIGIR 2020, Virtual Event, China, July 25-30, 2020*, Jimmy X. Huang, Yi Chang, Xueqi Cheng, Jaap Kamps, Vanessa Murdock, Ji-Rong Wen, and Yiqun Liu (Eds.). ACM, 639–648. doi:10.1145/3397271.3401063
- [9] Yinxuan Huang, Ke Liang, Zhuofan Dong, Xiaodong Qu, Tianxiang Wang, Yue Han, Jingao Xu, Bin Zhou, and Ye Wang. 2025. Flow Matching for Denoised Social Recommendation. In *Forty-second International Conference on Machine Learning, ICML 2025, Vancouver, BC, Canada, July 13-19, 2025*. OpenReview.net. <https://openreview.net/forum?id=gqKz6uCb11>
- [10] Mohsen Jamali and Martin Ester. 2010. A matrix factorization technique with trust propagation for recommendation in social networks. In *Proceedings of the 2010 ACM Conference on Recommender Systems, RecSys 2010, Barcelona, Spain, September 26-30, 2010*, Xavier Amatriain, Marc Torrens, Paul Resnick, and Markus Zanker (Eds.). ACM, Barcelona, Spain, 135–142. doi:10.1145/1864708.1864736
- [11] Kalervo Järvelin and Jaana Kekäläinen. 2002. Cumulated gain-based evaluation of IR techniques. *ACM Trans. Inf. Syst.* 20, 4 (2002), 422–446. doi:10.1145/582415.582418
- [12] Wei Jiang, Xinyi Gao, Guandong Xu, Tong Chen, and Hongzhi Yin. 2024. Challenging Low Homophily in Social Recommendation. In *Proceedings of the ACM on Web Conference 2024, WWW 2024, Singapore, May 13-17, 2024*, Tat-Seng Chua, Chong-Wah Ngo, Ravi Kumar, Hady W. Lauw, and Roy Ka-Wei Lee (Eds.). ACM, 3476–3484. doi:10.1145/3589334.3645460
- [13] Xin Jin and Jiawei Han. 2010. K-Means Clustering. In *Encyclopedia of Machine Learning*, Claude Sammut and Geoffrey I. Webb (Eds.). Springer US, Boston, MA, 563–564. doi:10.1007/978-0-387-30164-8_425
- [14] Dongyang Li, Jianshan Sun, Chongming Gao, Fuli Feng, and Kun Yuan. 2025. Independent or Social Driven Decision? A Counterfactual Refinement Strategy for Graph-Based Social Recommendation. *ACM Trans. Inf. Syst.* 43, 3 (2025), 76:1–76:27. doi:10.1145/3717830
- [15] Jiuqiang Li and Hongjun Wang. 2024. Graph Diffusive Self-Supervised Learning for Social Recommendation. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR 2024, Washington DC, USA, July 14-18, 2024*, Grace Hui Yang, Hongning Wang, Sam Han, Claudia Hauff, Guido Zucco, and Yi Zhang (Eds.). ACM, 2442–2446. doi:10.1145/3626772.3657962
- [16] Nian Li, Chen Gao, Depeng Jin, and Qingmin Liao. 2023. Disentangled Modeling of Social Homophily and Influence for Social Recommendation. *IEEE Trans. Knowl. Data Eng.* 35, 6 (2023), 5738–5751. doi:10.1109/TKDE.2023.3185388
- [17] Zongwei Li, Lianghao Xia, and Chao Huang. 2024. RecDiff: Diffusion Model for Social Recommendation. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management, CIKM 2024, Boise, ID, USA, October 21-25, 2024*, Edoardo Serra and Francesca Spezzano (Eds.). ACM, 1346–1355. doi:10.1145/3627673.3679630
- [18] Jie Liao, Wei Zhou, Fengji Luo, Junhao Wen, Min Gao, Xiuhua Li, and Jun Zeng. 2022. SocialLGN: Light graph convolution network for social recommendation. *Inf. Sci.* 589 (2022), 595–607. doi:10.1016/j.ins.2022.01.001
- [19] Sitao Luan, Chenqing Hua, Qincheng Lu, Jiaqi Zhu, Mingde Zhao, Shuyuan Zhang, Xiao-Wen Chang, and Doina Precup. 2022. Revisiting Heterophily For Graph Neural Networks. In *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*, Sanmi Koyejo, S. Mohamed, A. Agarwal, Danielle Belgrave, K. Cho, and A. Oh (Eds.). New Orleans, LA, USA. http://papers.nips.cc/paper_files/paper/2022/hash/092359ce5cf60a80e882378944bf1be4-Abstract-Conference.html
- [20] Hao Ma, Haixuan Yang, Michael R. Lyu, and Irwin King. 2008. SoRec: social recommendation using probabilistic matrix factorization. In *Proceedings of the 17th ACM Conference on Information and Knowledge Management, CIKM 2008, Napa Valley, California, USA, October 26-30, 2008*, James G. Shanahan, Sihem Amer-Yahia, Ioana Manolescu, Yi Zhang, David A. Evans, Aleksander Kolcz, Key-Sun Choi, and Abdur Chowdhury (Eds.). ACM, Napa Valley, CA, USA, 931–940. doi:10.1145/1458082.1458205
- [21] Hao Ma, Dengyong Zhou, Chao Liu, Michael R. Lyu, and Irwin King. 2011. Recommender systems with social regularization. In *Proceedings of the Forth International Conference on Web Search and Web Data Mining, WSDM 2011, Hong Kong, China, February 9-12, 2011*, Irwin King, Wolfgang Nejdl, and Hang Li (Eds.). ACM, 287–296. doi:10.1145/1935826.1935877
- [22] Wenze Ma, Yuxian Wang, Yanmin Zhu, Zhaobo Wang, Mengyuan Jing, Xuhao Zhao, Jiadi Yu, and Feilong Tang. 2024. MADM: A Model-agnostic Denoising Module for Graph-based Social Recommendation. In *Proceedings of the 17th ACM International Conference on Web Search and Data Mining, WSDM 2024, Merida, Mexico, March 4-8, 2024*, Luz Angelica Caudillo-Mata, Silvio Lattanzi, Andrés Muñoz Medina, Leman Akoglu, Aristides Gionis, and Sergei Vassilvitskii (Eds.). ACM, 501–509. doi:10.1145/3616855.3635784
- [23] J. MacQueen. 1967. Some Methods for Classification and Analysis of Multivariate Observations. In *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*. University of California Press, Berkeley, CA, USA, 281–297. <https://api.semanticscholar.org/CorpusID:6278891>
- [24] Miller McPherson, Lynn Smith-Lovin, and James M. Cook. 2001. Birds of a Feather: Homophily in Social Networks. *Annual Review of Sociology* 27 (2001), 415–444.
- [25] Thanh Tam Nguyen, Thanh Toan Nguyen, Matthias Weidlich, Jun Jo, Quoc Viet Hung Nguyen, Hongzhi Yin, and Alan Wee-Chung Liew. 2025. Handling Low Homophily in Recommender Systems With Partitioned Graph Transformer. *IEEE Transactions on Knowledge and Data Engineering* 37, 1 (2025), 334–350. doi:10.1109/TKDE.2024.3485880
- [26] Yuhan Quan, Jingtao Ding, Chen Gao, Lingling Yi, Depeng Jin, and Yong Li. 2023. Robust Preference-Guided Denoising for Graph based Social Recommendation. In *Proceedings of the ACM Web Conference 2023, WWW 2023, Austin, TX, USA, 30 April 2023 - 4 May 2023*, Ying Ding, Jie Tang, Juan F. Sequeda, Lora Aroyo, Carlos Castillo, and Geert-Jan Houben (Eds.). ACM, 1097–1108. doi:10.1145/3543507.3583374
- [27] Xubin Ren, Lianghao Xia, Jiashu Zhao, Dawei Yin, and Chao Huang. 2023. Disentangled Contrastive Collaborative Filtering. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR 2023, Taipei, Taiwan, July 23-27, 2023*, Hsin-Hsi Chen, Wei-Jou (Edward) Duh, Hen-Hsen Huang, Makoto P. Kato, Josiane Mothe, and Barbara Poblete (Eds.). ACM, 1137–1146. doi:10.1145/3539618.3591665
- [28] Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. 2009. BPR: Bayesian Personalized Ranking from Implicit Feedback. In *UAI 2009, Proceedings of the Twenty-Fifth Conference on Uncertainty in Artificial Intelligence, Montreal, QC, Canada, June 18-21, 2009*, Jeff A. Bilmes and Andrew Y. Ng (Eds.). AUAI Press, 452–461. https://www.auai.org/uai2009/papers/UAI2009_0139_48141db02b9f0b02bc7158819ebfa2c7.pdf
- [29] Xiao Sha, Zhu Sun, and Jie Zhang. 2022. Disentangling Multi-Facet Social Relations for Recommendation. *IEEE Trans. Comput. Soc. Syst.* 9, 3 (2022), 867–878. doi:10.1109/TCSS.2021.3108794
- [30] Youchen Sun, Zhu Sun, Yingpeng Du, Jie Zhang, and Yew Soon Ong. 2024. Self-Supervised Denoising through Independent Cascade Graph Augmentation for Robust Social Recommendation. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, KDD 2024, Barcelona, Spain, August 25-29, 2024*, Ricardo Baeza-Yates and Francesco Bonchi (Eds.). ACM, 2806–2817. doi:10.1145/3637528.3671958
- [31] Youchen Sun, Zhu Sun, Yingpeng Du, Jie Zhang, and Yew Soon Ong. 2025. Model-Agnostic Social Network Refinement with Diffusion Models for Robust Social Recommendation. In *Proceedings of the ACM on Web Conference 2025, WWW 2025, Sydney, NSW, Australia, 28 April 2025 - 2 May 2025*, Guodong Long, Michale Blumstein, Yi Chang, Liane Lewin-Eytan, Zi Helen Huang, and Elad Yom-Tov (Eds.). ACM, 370–378. doi:10.1145/3696410.3714683
- [32] Jiliang Tang, Xia Hu, and Huan Liu. 2013. Social recommendation: a review. *Soc. Netw. Anal. Min.* 3, 4 (2013), 1113–1133. doi:10.1007/S13278-013-0141-9

- [33] Aäron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation Learning with Contrastive Predictive Coding. *CoRR* abs/1807.03748 (2018). arXiv:1807.03748 [cs.LG] <https://arxiv.org/abs/1807.03748>
- [34] Hao Wang, Defu Lian, Hanghang Tong, Qi Liu, Zhenya Huang, and Enhong Chen. 2022. HyperSoRec: Exploiting Hyperbolic User and Item Representations with Multiple Aspects for Social-aware Recommendation. *ACM Trans. Inf. Syst.* 40, 2 (2022), 24:1–24:28. doi:10.1145/3463913
- [35] Tianle Wang, Lianghao Xia, and Chao Huang. 2023. Denoised Self-Augmented Learning for Social Recommendation. In *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence, IJCAI 2023, 19th-25th August 2023, Macao, SAR, China*. ijcai.org, 2324–2331. doi:10.24963/IJCAI.2023/258
- [36] Xiang Wang, Xiangnan He, Meng Wang, Fuli Feng, and Tat-Seng Chua. 2019. Neural Graph Collaborative Filtering. In *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR 2019, Paris, France, July 21-25, 2019*, Benjamin Piwowarski, Max Chevalier, Éric Gaussier, Yoelle Maarek, Jian-Yun Nie, and Falk Scholer (Eds.). ACM, 165–174. doi:10.1145/3331184.3331267
- [37] Xin Wang, Steven C. H. Hoi, Martin Ester, Jiajun Bu, and Chun Chen. 2017. Learning Personalized Preference of Strong and Weak Ties for Social Recommendation. In *Proceedings of the 26th International Conference on World Wide Web, WWW 2017, Perth, Australia, April 3-7, 2017*, Rick Barrett, Rick Cummings, Eugene Agichtein, and Evgeniy Gabrilovich (Eds.). ACM, 1601–1610. doi:10.1145/3038912.3052556
- [38] Xiang Wang, Hongye Jin, An Zhang, Xiangnan He, Tong Xu, and Tat-Seng Chua. 2020. Disentangled Graph Collaborative Filtering. In *Proceedings of the 43rd International ACM SIGIR conference on research and development in Information Retrieval, SIGIR 2020, Virtual Event, China, July 25-30, 2020*, Jimmy X. Huang, Yi Chang, Xueqi Cheng, Jaap Kamps, Vanessa Murdock, Ji-Rong Wen, and Yiqun Liu (Eds.). ACM, 1001–1010. doi:10.1145/3397271.3401137
- [39] Yuxiang Wei, Yupeng Shi, Xiao Liu, Zhilong Ji, Yuan Gao, Zhongqin Wu, and Wangmeng Zuo. 2021. Orthogonal Jacobian Regularization for Unsupervised Disentanglement in Image Generation. In *2021 IEEE/CVF International Conference on Computer Vision, ICCV 2021, Montreal, QC, Canada, October 10-17, 2021*. IEEE, 6701–6710. doi:10.1109/ICCV48922.2021.00665
- [40] Jiancan Wu, Xiang Wang, Fuli Feng, Xiangnan He, Liang Chen, Jianxun Lian, and Xing Xie. 2021. Self-supervised Graph Learning for Recommendation. In *SIGIR '21: The 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, Virtual Event, Canada, July 11-15, 2021*, Fernando Diaz, Chirag Shah, Torsten Suel, Pablo Castells, Rosie Jones, and Tetsuya Sakai (Eds.). ACM, 726–735. doi:10.1145/3404835.3462862
- [41] Le Wu, Junwei Li, Peijie Sun, Richang Hong, Yong Ge, and Meng Wang. 2022. DiffNet++: A Neural Influence and Interest Diffusion Network for Social Recommendation. *IEEE Trans. Knowl. Data Eng.* 34, 10 (2022), 4753–4766. doi:10.1109/TKDE.2020.3048414
- [42] Le Wu, Peijie Sun, Yanjie Fu, Richang Hong, Xiting Wang, and Meng Wang. 2019. A Neural Influence Diffusion Model for Social Recommendation. In *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR 2019, Paris, France, July 21-25, 2019*, Benjamin Piwowarski, Max Chevalier, Éric Gaussier, Yoelle Maarek, Jian-Yun Nie, and Falk Scholer (Eds.). ACM, 235–244. doi:10.1145/3331184.3331214
- [43] Lianghao Xia, Yizhen Shao, Chao Huang, Yong Xu, Huance Xu, and Jian Pei. 2023. Disentangled Graph Social Recommendation. In *39th IEEE International Conference on Data Engineering, ICDE 2023, Anaheim, CA, USA, April 3-7, 2023*. IEEE, 2332–2344. doi:10.1109/ICDE55515.2023.00180
- [44] Bo Yang, Yu Lei, Jiming Liu, and Wenjie Li. 2017. Social Collaborative Filtering by Trust. *IEEE Trans. Pattern Anal. Mach. Intell.* 39, 8 (2017), 1633–1647. doi:10.1109/TPAMI.2016.2605085
- [45] Shixiao Yang, Zhida Qin, Enjun Du, Pengzhan Zhou, and Tianyu Huang. 2024. Dual Social View Enhanced Contrastive Learning for Social Recommendation. *IEEE Transactions on Computational Social Systems* (2024).
- [46] Yonghui Yang, Le Wu, Yuxin Liao, Zhuangzhuang He, Pengyang Shao, Richang Hong, and Meng Wang. 2025. Invariance Matters: Empowering Social Recommendation via Graph Invariant Learning. In *Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR 2025, Padua, Italy, July 13-18, 2025*, Nicola Ferro, Maria Maistro, Gabriella Pasi, Omar Alonso, Andrew Trotman, and Suzan Verberne (Eds.). ACM, 2038–2047. doi:10.1145/3726302.3730013
- [47] Yonghui Yang, Le Wu, Zihan Wang, Zhuangzhuang He, Richang Hong, and Meng Wang. 2024. Graph Bottlenecked Social Recommendation. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, KDD 2024, Barcelona, Spain, August 25-29, 2024*, Ricardo Baeza-Yates and Francesco Bonchi (Eds.). ACM, 3853–3862. doi:10.1145/3637528.3671807
- [48] Junliang Yu, Hongzhi Yin, Min Gao, Xin Xia, Xiangliang Zhang, and Nguyen Quoc Viet Hung. 2021. Socially-Aware Self-Supervised Tri-Training for Recommendation. In *KDD '21: The 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Virtual Event, Singapore, August 14-18, 2021*, Feida Zhu, Beng Chin Ooi, and Chunyan Miao (Eds.). ACM, 2084–2092. doi:10.1145/3447548.3467340
- [49] Junliang Yu, Hongzhi Yin, Jundong Li, Qinyong Wang, Nguyen Quoc Viet Hung, and Xiangliang Zhang. 2021. Self-Supervised Multi-Channel Hypergraph Convolutional Network for Social Recommendation. In *WWW '21: The Web Conference 2021, Virtual Event / Ljubljana, Slovenia, April 19-23, 2021*, Jure Leskovec, Marko Grobelnik, Marc Najork, Jie Tang, and Leila Zia (Eds.). ACM / IW3C2, 413–424. doi:10.1145/3442381.3449844
- [50] Junliang Yu, Hongzhi Yin, Xin Xia, Tong Chen, Lizhen Cui, and Quoc Viet Hung Nguyen. 2022. Are Graph Augmentations Necessary?: Simple Graph Contrastive Learning for Recommendation. In *SIGIR '22: The 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, Madrid, Spain, July 11 - 15, 2022*, Enrique Amigó, Pablo Castells, Julio Gonzalo, Ben Carterette, J. Shane Culpepper, and Gabriella Kazai (Eds.). ACM, 1294–1303. doi:10.1145/3477495.3531937
- [51] Junliang Yu, Hongzhi Yin, Xin Xia, Tong Chen, Jundong Li, and Zi Huang. 2023. Self-supervised learning for recommender systems: A survey. *IEEE Transactions on Knowledge and Data Engineering* (2023).
- [52] Yu Zheng, Chen Gao, Xiang Li, Xiangnan He, Yong Li, and Depeng Jin. 2021. Disentangling User Interest and Conformity for Recommendation with Causal Embedding. In *WWW '21: The Web Conference 2021, Virtual Event / Ljubljana, Slovenia, April 19-23, 2021*, Jure Leskovec, Marko Grobelnik, Marc Najork, Jie Tang, and Leila Zia (Eds.). ACM / IW3C2, 2980–2991. doi:10.1145/3442381.3449788
- [53] Jiong Zhu, Yujun Yan, Lingxiao Zhao, Mark Heimann, Leman Akoglu, and Danai Koutra. 2020. Beyond Homophily in Graph Neural Networks: Current Limitations and Effective Designs. In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*, Hugo Larochelle, Marc'Aurelio Ranzato, Raia Hadsell, Maria-Florina Balcan, and Hsuan-Tien Lin (Eds.). <https://proceedings.neurips.cc/paper/2020/hash/58ae23d878a47004366189884c2f8440-Abstract.html>

A Projection Linearity under Social Aggregation

Fix $\mathbf{a}_u^{(l)}$ and define the projection operators:

$$\mathbf{P}_a(\mathbf{x}) \triangleq \frac{\mathbf{a}_u^{(l)} \mathbf{a}_u^{(l)\top}}{\|\mathbf{a}_u^{(l)}\|_2^2 + \epsilon} \mathbf{x}, \quad \mathbf{P}_a^\perp(\mathbf{x}) \triangleq \mathbf{x} - \mathbf{P}_a(\mathbf{x}). \quad (\text{A.1})$$

Since $\mathbf{P}_a(\cdot)$ is linear, it commutes with the linear social aggregation $\mathbf{m}_{u,S}^{(l)} = \sum_v w_{uv} \mathbf{s}_v^{(l)}$ (Eq. (7), $w_{uv} = 1/d_u^S$):

$$\mathbf{P}_a(\mathbf{m}_{u,S}^{(l)}) = \mathbf{P}_a\left(\sum_v w_{uv} \mathbf{s}_v^{(l)}\right) = \sum_v w_{uv} \mathbf{P}_a(\mathbf{s}_v^{(l)}), \quad (\text{A.2})$$

and the same holds for $\mathbf{P}_a^\perp(\cdot)$:

$$\mathbf{P}_a^\perp(\mathbf{m}_{u,S}^{(l)}) = \sum_v w_{uv} \mathbf{P}_a^\perp(\mathbf{s}_v^{(l)}). \quad (\text{A.3})$$

B Mechanism Effectiveness Analysis

To assess the effect of the learned gate, we conduct an inference-time intervention experiment. We freeze the trained model parameters and re-run a full forward propagation with a modified injection rule applied at every layer, without any retraining. Specifically, we replace the deviation-admission term when forming the filtered social injection:

- **Base (ODSR):** $\mathbf{s}_{u,\text{clean}}^{(l)} = \mathbf{s}_{u,\parallel}^{(l)} + \mathbf{g}_u^{(l)} \odot \mathbf{s}_{u,\perp}^{(l)}$.
- **All-pass Gate:** $\mathbf{s}_{u,\text{clean}}^{(l)} = \mathbf{s}_{u,\parallel}^{(l)} + \mathbf{s}_{u,\perp}^{(l)}$, equivalently setting $\mathbf{g}_u^{(l)} = \mathbf{1}$ for all users and layers.

Table 5 reports the performance gap relative to **Base (ODSR)**. Consistently across datasets, removing deviation control (**All-pass Gate**) degrades ranking performance, supporting that regulating off-anchor deviation is necessary under ranking supervision.

C Robustness of Intent-Overlap Mixedness

Tables 6–7 summarize the robustness check under $K_c=20$ with five random seeds. Across all datasets, Mixed ties dominate in the real graphs, whereas the shuffled-label Null test substantially changes the composition and r^+ distribution, indicating that the observed mid-range dominance is not a trivial artifact.

Table 5: Performance gap, measured by $\Delta\text{NDCG}@20$, relative to ODSR. Strictly negative values indicate that removing deviation control hurts performance.

Variant	Epinions	Yelp2018	Douban-Book
All-pass Gate	-0.0055	-0.0052	-0.0141

Robustness of r^+ (statistics). Tables 6–7 report mean $\pm$ std over five random seeds under $K_c=20$ for all datasets.

Null test (shuffled labels). To construct the Null (shuffle) baseline in Tables 6–7, we keep the interaction graph, social graph, and the learned item embeddings fixed, and randomly shuffle the item-to-intent assignments (i.e., cluster labels) to break the semantic correspondence between items and intent groups. This preserves the marginal cluster-size profile while disrupting the overlap structure induced by user interactions, so that any remaining mid-range dominance would mainly reflect random coincidence rather than structured preference overlap. We report both Real and Null (shuffle) statistics to verify that the observed r^+ patterns are not a trivial artifact of clustering or degree effects. As an additional diagnostic, shuffling item cluster labels substantially alters the overlap structure, supporting that the observed mid-range dominance reflects non-random overlap rather than a trivial clustering artifact.

D Intent Extraction Details

This section provides implementation details for the intent construction procedure used in Section 2. We treat *intent groups* as item clusters induced from training interactions, and derive user-level intent sets C_u accordingly.

Interaction-only item embeddings. We learn item embeddings from the interaction graph only, without using social links, on the training split, and use them as inputs for clustering. Concretely, we train a LightGCN model on $\mathcal{G}_I$ with the same optimization objective and split as in the main experiments, and take its pooled item embedding (Eq. (18)) as $\mathbf{z}_i \in \mathbb{R}^D$.

K-means clustering and model selection. We run K-means with K_c clusters, set to $K_c=20$ by default, on $\{\mathbf{z}_i\}$. We use multiple random initializations and select the partition with the lowest within-cluster

Table 6: Robustness of edge composition (mean $\pm$ std over seeds).

Dataset	Real			Null (shuffle)		
	Mixed (%)	Noise (%)	Signal (%)	Mixed (%)	Noise (%)	Signal (%)
Douban-Book	75.24 $\pm$ 0.59	2.97 $\pm$ 0.10	21.78 $\pm$ 0.61	58.30 $\pm$ 2.20	0.68 $\pm$ 0.12	41.03 $\pm$ 2.21
Epinions	75.13 $\pm$ 2.15	1.06 $\pm$ 0.13	23.82 $\pm$ 2.22	57.03 $\pm$ 2.85	0.19 $\pm$ 0.05	42.78 $\pm$ 2.90
Yelp2018	65.80 $\pm$ 1.55	0.62 $\pm$ 0.27	33.58 $\pm$ 1.67	77.44 $\pm$ 1.23	0.00 $\pm$ 0.00	22.56 $\pm$ 1.23

Table 7: Robustness of r^+ statistics (mean $\pm$ std over seeds).

Dataset	Real				Null (shuffle)			
	mean	median	q10	q90	mean	median	q10	q90
Douban-Book	0.633 $\pm$ 0.009	0.661 $\pm$ 0.012	0.200 $\pm$ 0.000	1.000 $\pm$ 0.000	0.725 $\pm$ 0.010	0.770 $\pm$ 0.024	0.314 $\pm$ 0.023	1.000 $\pm$ 0.000
Epinions	0.627 $\pm$ 0.026	0.623 $\pm$ 0.054	0.281 $\pm$ 0.031	1.000 $\pm$ 0.000	0.697 $\pm$ 0.014	0.684 $\pm$ 0.033	0.316 $\pm$ 0.033	1.000 $\pm$ 0.000
Yelp2018	0.701 $\pm$ 0.012	0.736 $\pm$ 0.018	0.317 $\pm$ 0.020	1.000 $\pm$ 0.000	0.612 $\pm$ 0.008	0.574 $\pm$ 0.039	0.250 $\pm$ 0.000	1.000 $\pm$ 0.000

sum of squared errors (SSE):

$$\text{SSE} = \sum_{k=1}^{K_c} \sum_{i \in \mathcal{V}_k} \|\mathbf{z}_i - \boldsymbol{\mu}_k\|_2^2, \quad (\text{D.4})$$

where $\mathcal{V}_k$ is the set of items assigned to cluster k and $\boldsymbol{\mu}_k$ is the centroid.

User intent sets. Given the selected clustering, each cluster index corresponds to an intent group. For each user u , we define C_u as the set of intent groups that contain items the user interacted with in the training split. These sets are then used to compute the directed intent overlap $r_{u \leftarrow v}^+$ in Eq. (1).

Reproducibility notes. To reduce dependence on a single clustering initialization, we repeat K-means with multiple random initializations and always select the partition with the lowest SSE, following the procedure in Section 2. We further report robustness statistics as mean $\pm$ std over five random seeds in Appendix C.

E Rotation Construction for $\Delta\text{TopShare}$

This section details the random-rotation construction used for computing $\Delta\text{TopShare}$ in Eq. (4) in Section 2.

Deviation vectors. For each directed social tie, we obtain a deviation vector $\mathbf{d} \in \mathbb{R}^D$ from the anchor-orthogonal decomposition described in Section 2. We compute $\text{TopShare}(\mathbf{d})$ as the fraction of ℓ_2 energy captured by the top- $\lceil \rho D \rceil$ dimensions ranked by $|\mathbf{d}_k|$ (with $\rho = 0.3$).

Random orthogonal rotation. To construct $\mathbf{d}_{\text{rot}} = \mathbf{Q}\mathbf{d}_{\text{real}}$, we sample an orthogonal matrix $\mathbf{Q} \in \mathbb{R}^{D \times D}$ from the Haar measure, which preserves $\|\mathbf{d}\|_2$ while randomizing the coordinate axes. In practice, we draw a Gaussian random matrix $\mathbf{G} \in \mathbb{R}^{D \times D}$ with i.i.d. entries, perform QR decomposition $\mathbf{G} = \mathbf{Q}\mathbf{R}$, and apply a diagonal sign correction using $\mathbf{R}$ so that $\mathbf{Q}$ is Haar-distributed.

Rotation baseline in $\Delta\text{TopShare}$. We compute the difference for each tie:

$$\Delta\text{TopShare} = \text{TopShare}(\mathbf{d}_{\text{real}}) - \text{TopShare}(\mathbf{d}_{\text{rot}}). \quad (\text{E.5})$$

If the deviation component exhibits dimension-wise concentration in the learned embedding basis, rotating the axes should weaken this concentration on average. A positive $\Delta\text{TopShare}$ indicates that deviation energy is more concentrated in the learned basis than under a random rotation, supporting the dimension-wise concentration evidence in Sec. 2.2.