

Distribution-aware Re-representations for Multi-Scenario Recommendations

Qi Sun*

The University of Melbourne
Melbourne, Australia
sunqs@student.unimelb.edu.au

Yulin Xu*

University of California Irvine
Irvine, CA, USA
yulinx8@uci.edu

Zelin Wang*

Independent
Shenzhen, China
zelinwang@gmail.com

Xiaoyu Kang

Institute of Information Engineering,
Chinese Academy of Sciences
Beijing, China
kangxiaoyu@iie.ac.cn

Keyan Jin

Macao Polytechnic University
Macao, China
p2317001@mpu.edu.mo

Jiechao Gao[†]

Stanford University
Stanford, CA, USA
jiechao@stanford.edu

Abstract

Modern applications provided personalized recommendations across diverse scenarios, including the homepage, local pages, and live streams on platforms like TikTok. These scenarios exhibit varying user behavior patterns, resulting in heterogeneous yet interrelated distributions. Existing Multi-Scenario Recommendation (MSR) methods usually use parameter-sharing networks for shared features and scenario-specific networks for unique features. However, these methods fail to handle different distribution across scenarios, resulting in representation entanglement and localization, which hinder effective knowledge transfer and compromise performance. In this paper, we propose a Distribution-aware Re-representations (DAR) method for MSR. Its core idea is to construct distribution-aware prototype spaces and learn disentangled re-representations around global prototypes. Specifically, DAR employs a Multi-gate Mixture of Experts (MMoE) to obtain scenario-shared representations, and uses independent networks to learn scenario-specific representations. These representations are then projected into scenario-shared and scenario-specific prototype spaces, producing scenario-shared re-representations (capturing global information) and scenario-specific re-representations (focusing on distributional differences). During this process, DAR utilizes Unbalanced Optimal Transport (UOT) to compute the transport relationships between representations and global prototypes, taking these as pseudo-labels for re-representation learning. Moreover, to prevent prototype entanglement, a matrix orthogonalization constraint ensures independence among global prototypes. The effectiveness of DAR is demonstrated through extensive offline experiments conducted on four datasets, as well as online A/B tests on a video platform.

CCS Concepts

• Information systems → Recommender systems.

* Equal contribution. [†] Corresponding author.



This work is licensed under a Creative Commons Attribution 4.0 International License.
SIGIR '26, Melbourne, VIC, Australia
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2599-9/2026/07
<https://doi.org/10.1145/3805712.3809554>

Keywords

Multi-Scenario Recommendation, Prototype learning, Optimal Transport

ACM Reference Format:

Qi Sun*, Yulin Xu*, Zelin Wang*, Xiaoyu Kang, Keyan Jin, and Jiechao Gao[†]. 2026. Distribution-aware Re-representations for Multi-Scenario Recommendations. In *Proceedings of the 49th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '26)*, July 20–24, 2026, Melbourne, VIC, Australia. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/3805712.3809554>

1 Introduction

Personalized recommendation has become an essential tool for addressing information overload, empowering users to efficiently discover content that aligns with their individual preferences [17, 29, 30, 34]. As user expectations and demands evolve, the rapid advancement of the Internet has further fueled the expansion of recommendation scenarios [6, 15, 26, 35, 42]. Modern platforms, such as Douyin, now provide personalized recommendation services across diverse and highly specialized scenarios [7, 16, 37], including the homepage, local pages, and live streams, as illustrated in Fig. 1.

Each scenario has different user behavior patterns, resulting in heterogeneous but interrelated data distributions [22, 28, 39]. Building separate models for each scenario may capture scenario-specific characteristics but fails to leverage shared user behaviors and common interests across scenarios [22, 39]. This approach also increases maintenance costs for enterprises [1, 22]. Therefore, effectively integrating multi-scenario data and improving model efficiency have become key challenges, making Multi-Scenario Recommendation (MSR) a critical topic [10, 21, 36].

Existing Multi-Scenario Recommendation (MSR) methods have adopted various strategies to address shared and scenario-specific challenges. Traditional methods like ShareBottom [2] and expert-based models such as MMoE [22] and M3OE [40] dynamically allocate knowledge but often fail to handle distributional differences across scenarios. Approaches like PLE [28] and AdaSparse [38] focus on feature separation, while attention-based models like STAR [26] and meta-learning methods such as M2M [39] enhance scenario adaptation. However, these methods fail to handle the distributional differences across scenarios, leading to representation

entanglement and localization, which hinder effective knowledge transfer and compromise performance.

To address these challenges, we propose **Distribution-aware Re-representations (DAR)**, a novel method for MSR that constructs distribution-aware prototype spaces and learns disentangled re-representations around global prototypes to effectively solve the issue of representation entanglement and localization. Figure 2 illustrates the core idea of DAR. DAR employs a Multi-gate Mixture of Experts (MMoE) model [22] to capture scenario-shared representations across scenarios while leveraging independent networks to learn scenario-specific representations for different scenarios. These representations are subsequently projected into scenario-shared and scenario-specific prototype spaces, generating scenario-shared re-representations that capture global information and scenario-specific re-representations that emphasize distributional differences. To ensure precise projection into each prototype space, we incorporate Unbalanced Optimal Transport (UOT) [24] into the learning process, given the inherent imbalance between individual representations and the overall distribution represented by prototypes. UOT computes the transport relationships between representations and global prototypes, providing adaptive alignment and serving as pseudo-labels to guide the learning of disentangled re-representations. These re-representations are then concatenated and fed into a classifier to get the final predictions. Additionally, to prevent prototype entanglement, a matrix orthogonalization constraint is imposed to ensure the independence of global prototypes. Our contributions can be summarized as follows:

- We propose a Distribution-aware Re-representation (DAR) framework, addressing representation entanglement and localization for Multi-Scenario Recommendation (MSR).
- We introduce a novel re-representation scheme based on Unbalanced Optimal Transport (UOT), enabling effective alignment and disentanglement of scenario-shared and scenario-specific features.
- We demonstrate the effectiveness of DAR through comprehensive offline experiments and online A/B testing, showcasing its superiority in improving recommendation performance.

2 Related Work

In recent times, there has been a noticeable shift towards Multi-Scenario Recommendation (MSR) tasks, driven by the rapid expansion of user bases and web content. Online platforms are increasingly segmenting user groups and content themes into distinct scenarios based on varied attributes. This strategy aligns closely with the principles of multi-task learning, prompting researchers to investigate domain-transfer technologies to address these challenges. Traditional methods such as ShareBottom [2] provide foundational insights by leveraging shared networks to capture common patterns while using scenario-specific towers for unique feature modeling. However, these methods often struggle to handle complex inter-scenario relationships and lack flexibility in dynamic environments. To address these limitations, more advanced architectures have been proposed to enhance inter-scenario modeling and adaptability. One prominent method involves using Mixture-of-Experts



Figure 1: An example of MSR in Douyin. Different pages, such as the homepage, local page, and live page, correspond to distinct scenarios.

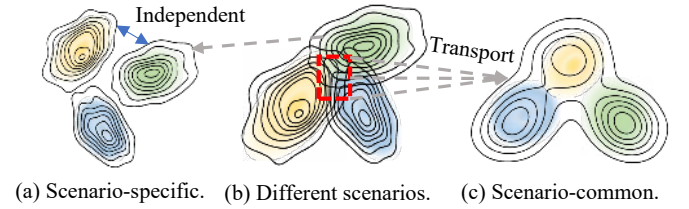


Figure 2: The core idea of our proposed method.

(MoE) architectures to manage the diversity across different domains [18, 31, 43]. For instance, the Mario model [31] is designed to capture domain-specific information through feature scaling modules, dynamically adjusting to domain signals using a MoE structure. This model excels in adapting to various domain requirements by leveraging a mixture of specialized expert models. Additionally, HiNet [43] employs a hierarchical structure that constructs multiple layers for information extraction. This approach ensures that while broad domain information is effectively captured, the specific features of each domain are preserved. HiNet’s multi-layered framework allows it to balance generalization and specialization, making it highly effective in complex recommendation scenarios. Further innovative methodologies have also been developed. PEPNet [3] processes bottom-level inputs using gating units and integrates EPNet [9] for domain feature selection. This method is complemented by PPNET [13], which combines multi-task information to enhance the overall recommendation system’s performance. M3OE [40] extends this idea by introducing a framework with multi-level MoE modules, capturing shared, domain-specific, and task-specific attributes. By integrating multiple layers of domain-specific and multi-task learning capabilities, these models aim to provide more accurate and personalized recommendations. In recent years, researchers have adopted various innovative approaches to tackle the challenges in domain modeling and adaptation. One such approach

involves leveraging dual-structured models like STAR [26], which combines a domain-specific tower with a domain-shared tower to effectively capture and enhance both unique and shared information, thereby improving performance across domains. Attention mechanisms have also been pivotal in advancing domain adaptation. Models such as SAR-Net [25], SAML [4], and M2M [39] utilize attention layers to facilitate seamless knowledge transfer across different domains, resulting in significant performance improvements. Meanwhile, HAMUR [20] focuses on enhancing distribution adaptation through domain adapters, which leads to better prediction outcomes across diverse domains. Prompt-based technologies have found their place in domain adaptation strategies as well. PLATE [33], for instance, employs widely-used NLP prompt techniques to boost adaptation performance. Additionally, research efforts have explored domain knowledge transfer through embedding alignment [23]. For instance, CausalInt [32] refines domain recommendations using causal inference techniques, while AdaSparse [38] introduces novel pruning strategies across domains to enhance domain adaptation.

3 preliminary

Optimal Transport (OT) [24] is a mathematical framework designed to transport mass between two distributions with minimal cost. A convenient way to formalize OT is by working within the probability simplices. Let $\Sigma_r := x \in \mathbb{R}_+^r : x^\top \mathbf{1}_r = 1$ denote the r -dimensional probability simplex, where $\mathbf{1}_r$ is a vector of ones. For two probability distributions $\alpha \in \Sigma_r$ and $\beta \in \Sigma_c$, the transport polytope $U(\alpha, \beta)$ is defined as:

$$U(\alpha, \beta) = \left\{ Q \in \mathbb{R}_+^{r \times c} \mid Q \mathbf{1}_c = \alpha, Q^\top \mathbf{1}_r = \beta \right\}. \quad (1)$$

Here, $Q \in \mathbb{R}_+^{r \times c}$ is the transport matrix, where Q_{ij} represents the mass transported from the i -th element of α to the j -th element of β . The constraints ensure marginal distributions match α and β . To solve OT, we use a similarity matrix $M \in \mathbb{R}^{r \times c}$, often constructed as the negative of a distance matrix or as a similarity measure (e.g., cosine similarity). The goal is to find Q maximizing $\text{Tr}(Q^\top M)$:

$$\text{OT}(M, \alpha, \beta) = \arg \max_{Q \in U(\alpha, \beta)} \left(\text{Tr}(Q^\top M) \right), \quad (2)$$

To avoid sparse transport plans, an entropy regularization term $H(Q) = -\sum_{i,j} Q_{ij} \log Q_{ij}$ is introduced, yielding:

$$\text{OT}^\varepsilon(M, \alpha, \beta) = \arg \max_{Q \in U(\alpha, \beta)} \left(\text{Tr}(Q^\top M) + \varepsilon H(Q) \right), \quad (3)$$

where $\varepsilon > 0$ is the entropic regularization parameter that controls the balance between maximizing similarity and promoting smoothness. This problem can be efficiently solved using Sinkhorn's algorithm [8], which yields the optimal transport plan in the form:

$$Q^* = \text{Diag}(\mathbf{u}) \exp(M/\varepsilon) \text{Diag}(\mathbf{v}), \quad (4)$$

where $\mathbf{u}$ and $\mathbf{v}$ are nonnegative scaling vectors determined by iterative normalization of rows and columns, ensuring the marginal constraints are satisfied. However, classical OT assumes exact mass conservation, making it unsuitable for partial transport scenarios where the source and target distributions do not have the same total mass. To address this limitation, Unbalanced Optimal Transport (UOT) [5] relaxes the marginal constraints by introducing penalties

based on the Kullback-Leibler (KL) divergence. The UOT objective is formulated as:

$$\text{UOT}^\kappa(M, \alpha, \beta) = \arg \max_{Q \geq 0} \left[\text{Tr}(Q^\top M) + \varepsilon H(Q) - \kappa \left(D_{\text{KL}}(Q \mathbf{1}_c \parallel \alpha) + D_{\text{KL}}(Q^\top \mathbf{1}_r \parallel \beta) \right) \right]. \quad (5)$$

Here, the parameter $\kappa > 0$ controls the penalty strength for deviating from the exact marginal constraints. The KL divergence $D_{\text{KL}}(\mathbf{p} \parallel \mathbf{q}) = \sum_i p_i \log(p_i/q_i)$ measures the discrepancy between the marginal distributions $Q \mathbf{1}_c$ and α (similarly for $Q^\top \mathbf{1}_r$ and β). A larger κ enforces stricter adherence to the original marginals, while a smaller κ allows more flexibility in partial transport.

The UOT problem can be efficiently solved using a generalized Sinkhorn algorithm, which extends the classical entropic OT approach while enabling greater flexibility in scenarios where exact mass conservation is not required.

4 METHODOLOGY

In this section, we introduce the proposed method in detail, which is mainly composed of six parts. Section 4.1 provides the formal definition of the Multi-Scenario Recommendation (MSR) problem. Section 4.2 discusses the basic representation learning framework. Section 4.3 describes the construction of global prototypes. Section 4.4 explains the re-representation learning process with Unbalanced Optimal Transport (UOT). Section 4.5 presents the loss functions used for optimization and Section 4.6 analyzes the time and space complexity of the proposed method.

4.1 Problem Definition

Let $\mathcal{U} = \{u_1, \dots, u_{|\mathcal{U}|}\}$, $\mathcal{I} = \{i_1, \dots, i_{|\mathcal{I}|}\}$ and $\mathcal{S} = \{1, 2, \dots, S\}$ denote the set of users, items and scenarios, respectively, where $|\mathcal{U}|$ is the number of users, $|\mathcal{I}|$ is the number of items, and S is the number of scenarios. The user-item historical interactions are represented by $\mathcal{D} = \{(u, i, s, y) \mid u \in \mathcal{U}, i \in \mathcal{I}, s \in \mathcal{S}\}$, where $y \in \{0, 1\}$ denotes the binary label. The primary goal of Multi-Scenario Recommendation is to develop a robust scoring function $F(\cdot \mid \Theta)$, trained on the dataset $\mathcal{D}$, that can accurately predict a user's preference for a given item across various scenarios. Here, Θ represents the set of parameters that define the scoring function F , which determines the likelihood of user u being interested in item i within the specific context of scenario s .

4.2 Basic Representation Learning

In this section, we describe the basic representation learning framework, which provides the foundational representations for subsequent models in the proposed method. It consists of two components: scenario-shared representation and scenario-specific representation learning. These two types of representations serve as the basis for constructing more refined features in later stages.

4.2.1 Scenario-shared Representation. We first utilize a shared feature extraction network to capture shared representations across multiple scenarios. This is based on the Multi-gate Mixture of Experts (MMoE) architecture [22], which allows for flexible parameter

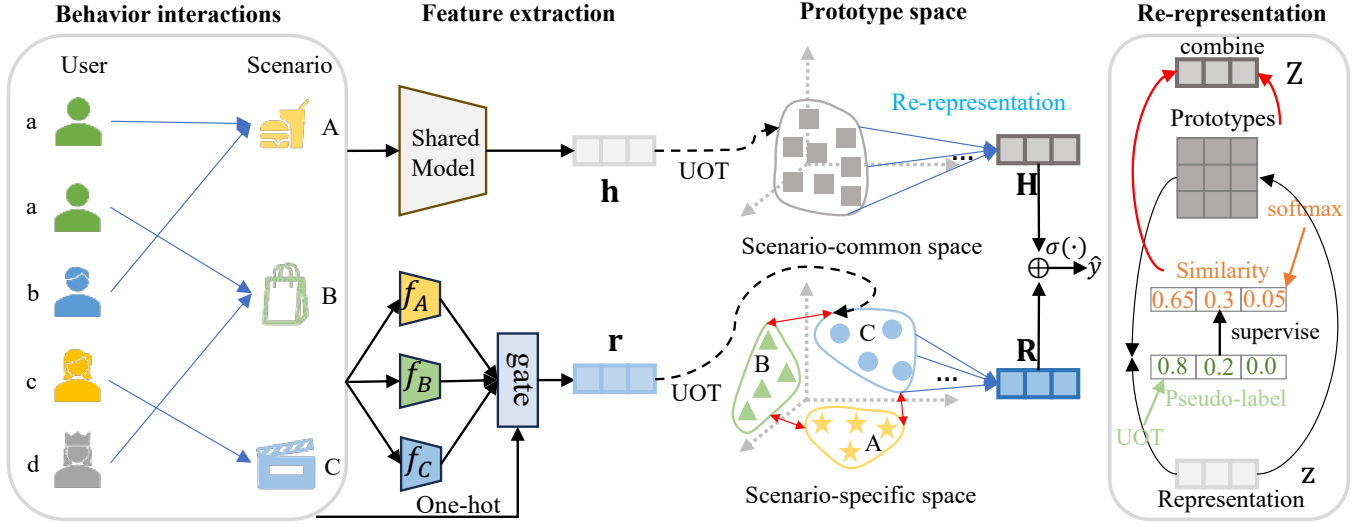


Figure 3: The architecture of DAR.

sharing and consistency of different scenarios. The shared representation of i -th sample is:

$$\mathbf{h}_i = \sum_{m=1}^M g_m^s(x_i) * Q_m(x_i), \quad (6)$$

where $Q_m(x)$ denotes the m -th expert network, which is composed of multi-layer perceptron (MLP) with ReLU activation function. Here, M represents the total number of expert networks. And $g^s(x)$ denotes the gating network for scenario s , which is a simply linear transformation of input x_i with a softmax layer:

$$g^s(x_i) = \text{softmax}(W_g x_i), \quad (7)$$

where W_g is the weight matrix for the gating network. The use of the MMoE [22] architecture ensures that each scenario benefits from shared knowledge.

4.2.2 Scenario-specific Representation. To capture unique characteristics of each scenario, we employ scenario-specific networks, which model the distinct distributional patterns and user behaviors within individual scenarios. For the i -th sample in scenario s , the scenario-specific representation is computed as:

$$\mathbf{r}_i = f^s(x_i), \quad (8)$$

where $f^s(\cdot)$ represents the scenario-specific network, implemented as a multi-layer perceptron (MLP) with ReLU activations. The scenario-specific networks are independently parameterized, ensuring that each network learns scenario-specific variations without interference from other scenarios.

4.3 Global Prototypes Construction

To construct global prototypes, we first train the recommendation model $F(\cdot; \Theta)$ using the basic representations obtained from the shared and scenario-specific networks. The scoring function $F(\cdot; \Theta)$ is parameterized as:

$$\hat{y}_i = F([\mathbf{h}_i; \mathbf{r}_i]; \Theta), \quad (9)$$

where $\mathbf{h}_i$ and $\mathbf{r}_i$ denote the scenario-shared and scenario-specific representations, respectively. The model is optimized using the binary cross-entropy loss:

$$\mathcal{L}_{\text{pretrain}} = -\frac{1}{|\mathcal{D}|} \sum_{(u,i,y) \in \mathcal{D}} [y \log(\hat{y}) + (1-y) \log(1-\hat{y})]. \quad (10)$$

After training the recommendation model $F(\cdot; \Theta)$ with the basic representations, we construct prototype spaces by applying K-Means clustering to the two type representations. The clustering process can be unified as:

$$\min_{\mathbf{P}} \frac{1}{|\mathcal{D}|} \sum_{i=1}^I \mathcal{D} \left| \min_{\mathbf{m}_i \in \{0,1,\dots,|\mathcal{S}|\}^K} \|\mathbf{z}_i - \mathbf{m}_i^T \mathbf{P}\| \right|, \quad \text{s.t.}; \mathbf{m}_i^T \mathbf{1}_K = 1, \quad (11)$$

where $\mathbf{z}_i$ represents the input representation for clustering, $\mathbf{P} \in \mathbb{R}^{K \times d}$ is the centroid matrix (prototypes), and $\mathbf{y}_i$ denotes the cluster assignment for the i -th representation.

For the shared representations $\mathbf{h}_i$, we apply K-Means clustering across all samples to generate K global prototypes $\mathbf{P}_{\text{share}}$. These shared prototypes capture global characteristics that are consistent across scenarios. For the scenario-specific representations $\mathbf{r}_i$, we cluster the samples within each scenario $s \in \mathcal{S}$. This generates $|\mathcal{S}|$ local prototypes $\mathbf{P}_s$ for each scenario, where $|\mathcal{S}|$ is the number of clusters for scenario s . These prototypes focus on the unique characteristics of each scenario. The use of K-Means for initializing global prototypes offers two main benefits: (i) it preserves the global data structure derived from the trained $F(\cdot; \Theta)$, even in noisy conditions, and (ii) it accelerates the convergence of the second training process compared to random initialization.

4.4 Re-representation Learning with UOT

To refine the representations and ensure alignment with the prototype spaces, we employ Unbalanced Optimal Transport (UOT), which generates pseudo-labels to supervise the re-representation

learning process. This is applied to both scenario-shared and scenario-specific representations $\mathbf{h}_i$ and $\mathbf{r}_i$, denoted by $\mathbf{z}_i$ for simplification.

In our setting, re-representation involves associating individual representations $\mathbf{z}_i \in \mathbb{R}^d$ with a set of prototypes $\mathbf{P} = \{\mathbf{p}_k\} \in \mathbb{R}^{K \times d}$. However, a single representation $\mathbf{z}_i$ only corresponds to a small subset of the prototypes, leading to an inherent imbalance in the re-representation process. UOT is particularly suited to handle such situations, as it allows for partial transport between a single sample and multiple prototypes, effectively addressing the imbalance and ensuring meaningful supervision.

4.4.1 Pseudo-label generations with UOT. We begin by defining a cost matrix $\mathbf{M}$ that quantifies the dissimilarity between each representation $\mathbf{z}_i$ and prototype $\mathbf{p}_k$ based on cosine similarity:

$$\mathbf{M}_{i,k} = 1 - \cos(\mathbf{z}_i, \mathbf{p}_k), \quad (12)$$

where $\cos(\cdot, \cdot)$ computes the cosine similarity. A smaller $\mathbf{M}_{i,k}$ indicates higher similarity between $\mathbf{z}_i$ and $\mathbf{p}_k$.

Using the cost matrix $\mathbf{M}$, we apply UOT to compute the optimal transport plan $\mathbf{Q}^*$:

$$\mathbf{Q}^* = \text{UOT}^K(\mathbf{M}, \boldsymbol{\alpha}, \boldsymbol{\beta}), \quad (13)$$

where $\boldsymbol{\alpha} = \delta_{\mathbf{z}_i}$ is a Dirac delta distribution representing the individual sample $\mathbf{z}_i$, and $\boldsymbol{\beta} = \sum_{k=1}^K \delta_{\mathbf{p}_k}$ is the weight distribution over the K prototypes $\mathbf{p}_k$. The regularization parameter κ controls the entropic smoothing of the transport plan. UOT allows for partial transport, ensuring that each representation $\mathbf{z}_i$ is softly associated with multiple prototypes based on their similarities.

4.4.2 Re-representation learning with pseudo-label. The transport plan $\mathbf{Q}^*$ acts as pseudo-labels to guide the re-representation learning process. Specifically, we refine the representations by minimizing the following supervised loss:

$$\mathcal{L}_{super}^{(i)}(\mathbf{z}_i, \mathbf{P}) = -\frac{1}{K} \sum_{k=1}^K \mathbf{Q}_{i,k}^* \log \frac{\exp(\mathbf{z}_i^T * \mathbf{p}_k / \tau)}{\sum_{j=1}^K \exp(\mathbf{z}_i^T * \mathbf{p}_j / \tau)}, \quad (14)$$

where τ is a temperature parameter controlling the sharpness of the distribution. This loss ensures that each $\mathbf{z}_i$ is re-represented in a way that reflects its relationships with the prototypes $\mathbf{p}_k$. The supervised loss essentially pushes each $\mathbf{z}_i$ closer to the prototypes it is most similar to (as indicated by $\mathbf{Q}^*$), while simultaneously discouraging similarity to unrelated prototypes. The use of UOT ensures that this re-representation process is robust to the imbalance between individual representations and the prototype space.

Finally, the re-representation $\mathbf{Z}_i$ are updated by combining the learned representations with their prototype assignments:

$$\mathbf{Z}_i = \sum_{k=1}^K \frac{\exp(\mathbf{z}_i^T * \mathbf{p}_k)}{\sum_{j=1}^K \exp(\mathbf{z}_i^T * \mathbf{p}_j)} * \mathbf{p}_k. \quad (15)$$

4.4.3 Usage in DAR. In the proposed DAR framework, the re-representation process is applied twice to handle both the scenario-shared and scenario-specific representations. First, the scenario-shared representations $\mathbf{h}_i$ are re-represented into $\mathbf{H}_i$ using the scenario-shared prototype space $\mathbf{P}_{share}$. This step captures global characteristics across all scenarios, ensuring that the shared representations align with the common semantic structure. Second, the scenario-specific representations $\mathbf{r}_i$ are re-represented into $\mathbf{R}_i$

using the scenario-specific prototype space $\mathbf{P}_s$ s -th scenario. This step refines the scenario-specific features by embedding them into a structured space tailored to each scenario's unique distribution. Therefore, for a sample i belonging to scenario s , the loss for its re-representation learning is:

$$\mathcal{L}_{usage} = \frac{1}{|\mathcal{D}|} \sum_{i=1}^{|\mathcal{D}|} \mathcal{L}_{super}^{(i)}(\mathbf{h}_i, \mathbf{P}_{share}) + \mathcal{L}_{super}^{(i)}(\mathbf{r}_i, \mathbf{P}_s). \quad (16)$$

4.5 Loss Functions

4.5.1 Inter-scenario orthogonalization loss. To ensure orthogonalization between prototypes across different scenarios, we propose a regularization term based on the similarity of prototype centers. Let the prototype matrix for scenario s be $\mathbf{P}_s = \{\mathbf{p}_{s,1}, \mathbf{p}_{s,2}, \dots, \mathbf{p}_{s,K}\}$, where each prototype is a d -dimensional vector. The prototype center for scenario s is defined as:

$$\bar{\mathbf{p}}_s = \frac{1}{K} \sum_{i=1}^K \mathbf{p}_{s,i}. \quad (17)$$

To encourage orthogonalization between scenarios, we minimize the squared cosine similarity between the prototype centers of different scenarios:

$$\mathcal{L}_{orth} = \sum_{1 \leq s < t \leq S} \left(\frac{\bar{\mathbf{p}}_s^T \bar{\mathbf{p}}_t}{\|\bar{\mathbf{p}}_s\| \|\bar{\mathbf{p}}_t\|} \right)^2. \quad (18)$$

4.5.2 Total loss. Finally, we concatenate the scenario-shared re-representation $\mathbf{H}_i$ and the scenario-specific re-representation $\mathbf{R}_i$ of sample i , and input them into the classifier to obtain the final prediction value:

$$\hat{y} = \sigma(\mathbf{H}_i \oplus \mathbf{R}_i) \quad (19)$$

We use binary cross-entropy loss to measure the prediction error:

$$\mathcal{L}_{bce} = \frac{1}{|\mathcal{D}|} \sum_{i=1}^{|\mathcal{D}|} y_i \log \hat{y}_i + (1 - y_i) \log(1 - \hat{y}_i). \quad (20)$$

The total loss function incorporates both the binary cross-entropy loss and the optimal transport loss:

$$\mathcal{L}_{total} = \mathcal{L}_{bce} + \gamma_1 * \mathcal{L}_{usage} + \gamma_2 * \mathcal{L}_{orth}, \quad (21)$$

where γ_1 and γ_2 are hyper-parameters.

4.6 Complexity and Scalability Discussions

4.6.1 Time complexity. During training, the main time cost comes from the UOT loss, which aligns a single sample with a set of K prototypes. Using the Sinkhorn algorithm [8], the time complexity for computing the UOT loss for a single sample is $O(K \log K)$. The independence term $\mathcal{L}_{orth}$ introduces a small computational overhead. With $|\mathcal{S}|$ scenarios, the total number of pairs is $\binom{|\mathcal{S}|}{2} = \frac{|\mathcal{S}|(|\mathcal{S}|-1)}{2}$, resulting in time complexity of: $O(|\mathcal{S}|^2)$.

In the inference phase, UOT is not required. Instead, the model calculates similarities between the representations and prototypes. This results in a time complexity of $O(K)$ per sample.

Table 1: Statistics information of datasets.

Dataset	Scenario	#Interaction	#User	#Item
Douban	S1	227,251	2,212	95,872
	S2	179,847	1,820	79,878
	S3	1,278,401	2,712	34,893
KuaiRand	S1	2,407,352	961	1,596,491
	S2	7,760,237	991	2,741,383
	S3	895,385	171	332,210
	S4	402,366	832	547,908
	S5	183,403	832	43,106
Ali-CCP	S1	32,236,951	89,283	465,870
	S2	639,897	2,561	188,610
	S3	52,439,671	150,471	467,122
Industry	S1	64,475,979	50,000	1,365,660
	S2	1,588,512	10,000	152,601
	S3	54,277,815	100,000	791,826

4.6.2 Space complexity. The primary space overhead in DAR comes from storing the prototypes. For a scenario-shared prototype $\mathbf{P}_{share}$ and $|S|$ scenario-specific prototypes $\{\mathbf{P}_1, \dots, \mathbf{P}_{|S|}\}$, each with a feature dimension of d , the total space complexity is: $O((|S| + 1) * K * d)$.

4.6.3 Scalability. The DAR framework is designed to be scalable in both time and space complexity. In terms of time, UOT is executed only twice during training, regardless of the number of scenarios $|S|$, ensuring that the time complexity per sample remains $O(K \log K)$. During inference, UOT is not required, and the similarity computation with prototypes results in a per-sample complexity of $O(K)$, making it efficient for real-time applications. For space complexity, the memory cost is dominated by the storage of prototypes. Real-world applications typically involve a limited number of scenarios. Therefore, the overall space complexity remains practical and does not significantly increase as the dataset size grows. These properties ensure that DAR is both computationally efficient and scalable for large-scale multi-scenario recommendation tasks.

5 Experiment

5.1 Experimental Settings

5.1.1 Datasets. We evaluate the effectiveness of DAR on four recommendation datasets. The split ratio of training, validation, and test is 7:1:2.

- **KuaiRand** [19]: The KuaiRand dataset provides an unbiased benchmark for Scenario-Wise Rec, derived from random video exposures on the Kuaishou App. It includes 11 million interactions across 1,000 users and 4 million videos. To simplify evaluation, data from the top five scenarios, representing distinct advertising positions, were selected.
- **Douban** [41]: This dataset, derived from the Douban platform, comprises three subsets: book, music, and movie, sharing a common user base and treated as separate scenarios. Following prior research, ratings above 3 are classified as positive, and those of 3 or below as negative.
- **Ali-CCP** [19]: Ali-CCP is a large-scale dataset for CTR prediction. It was collected from the traffic logs of Taobao's recommendation system. In this dataset, the context feature

"301" serves as a scenario identifier, indicating the position from which the interaction sample originates.

- **Industry:** We collected user interaction history data from three scenarios on a video platform: the homepage, the local city page, and the following page, and set finish playing as the task.

6 Baselines

We use the following state-of-the-art MSR/MTR model as baselines:

- **ShareBottom** [2]: It uses a shared network to learn common representations across scenarios, capturing shared patterns and information. Separate network towers are then applied for scenario-specific modeling.
- **MMOE** [22]: The Multi-gate Mixture-of-Experts (MMOE) model is a multi-task learning framework that integrates expert networks to handle diverse tasks. It uses gating networks to dynamically distribute task-specific knowledge, improving adaptability in complex scenarios.
- **PLE** [28]: The Progressive Layered Extraction (PLE) model separates shared and task-specific components to tackle negative transfer and task correlations in multi-task learning, using progressive routing to extract deeper features.
- **STAR** [26]: The Star Topology Adaptive Recommender (STAR) model integrates a shared network for common features and scenario-specific networks tailored to each scenario.
- **SAR-Net** [25]: The Scenario-Aware Ranking Network uses attention modules and scenario-specific expert networks with a multi-scenario gating module to handle biased logs.
- **M2M** [39]: The Multi-Scenario Multi-Task Meta-Learning model uses a meta unit for inter-scenario relationships, meta attention for task-specific correlations, and a meta tower to enhance scenario-specific features.
- **AdaSparse** [38]: It adaptively learns sparse structures in scenario-specific models. It uses a pruner network for scenario-level pruning, combining binary and scale-based strategies to remove redundant information.
- **CausalInt** [32]: CausalInt leverages causal inference to address biases in multi-scenario recommendation. By modeling causal relationships, it disentangles spurious correlations from meaningful patterns, improving the reliability of knowledge transfer across scenarios and enhancing recommendation accuracy in diverse contexts.
- **PPNet & EPNet** [3]: PPNet and EPNet, part of the Parameter and Embedding Personalized Network (PEPNet), handle multi-task recommendations under multi-scenario settings. EPNet fuses features with different importance for users, while PPNet modifies parameters for different tasks.
- **M3OE** [40]: It introduces a framework consisting of three Mixture-of-Experts (MoE) modules to learn common, domain-specific, and task-specific attributes, along with a two-level fusion mechanism that enables precise control over feature extraction and fusion across different domains and tasks.

6.0.1 Evaluation Metrics. To evaluate the performance of proposed model, we adopt two commonly used metrics in MSR tasks:

- **AUC** (Area Under the ROC Curve) [12] measures the ability of model to distinguish between positive and negative

classes. A higher AUC value indicates better discrimination capability.

- **LogLoss** (Logarithmic Loss) [11] quantifies the accuracy of predicted probabilities. It penalizes incorrect predictions, with lower LogLoss values indicating better performance.

6.0.2 Hyper-parameter and training details. The input embedding dimension is set to 32 for all methods. In shared feature extraction network, the $Q(\cdot)$ is a three-layer MLP with activation function of ReLU, and M is set to 3. The number of prototype K is searched in $\{8, 16, 32, 64\}$. The hyper-parameter of training, γ_1 and γ_2 , which control the degree of usage loss and orthogonalization loss, respectively, are both searched in the range of $\{0.1, 0.2, \dots, 1.0\}$. We optimize all models with Adam [14] optimizer with batch sizes of $\{256, 512, 1024\}$. We use grid search to find the optimal hyperparameters. The learning rate is searched in $\{1e-3, 1e-4, 1e-5\}$. To avoid overfitting, we set the dropout [27] to 0.5 and the patience of earlystop to 5 epochs.

6.1 Overall Performance

We compare DAR with several baseline models on four datasets. From Table 2, we can safely draw the following conclusions:

(1) Simple parameter-sharing models, such as ShareBottom and MMoE, struggle to adapt to the diverse distributions present in multi-scenario datasets. These methods rely on rigid shared representations, which leads to suboptimal performance when faced with significant differences between scenarios. As a result, their average AUC lags behind multi-scenario recommendation methods by approximately 2.1%, highlighting the inherent limitations of parameter-sharing strategies in scenarios requiring both global knowledge transfer and local adaptation. (2) In contrast, multi-scenario recommendation methods such as PLE, SAR-Net, and STAR introduce mechanisms to balance scenario-specific feature learning with shared knowledge transfer. Among these, STAR demonstrates strong performance across datasets due to its domain-specific representation capabilities, achieving a notable improvement over simple parameter-sharing methods. However, in datasets with significant distribution imbalances, such as Douban, STAR falls short compared to alignment-based models, with an average gap of 1.2% in AUC, reflecting the challenges of handling diverse distributions solely through representation learning. (3) Feature alignment models, such as M2M and CausalInt, address these challenges by aligning feature distributions across scenarios, resulting in competitive LogLoss values across datasets with high scenario diversity, such as KuaiRand. However, their reliance on global alignment often leads to a lack of granularity in capturing scenario-specific nuances, limiting their effectiveness in datasets like Douban, where scenario-specific differences are more pronounced. This trade-off underscores the need for a more integrated approach that balances global alignment with local customization. (4) Enhanced expert models, represented by PPNet and EPNet, focus on refining expert network designs through dynamic weighting mechanisms, achieving moderate improvements in datasets with balanced or moderately diverse scenarios. For example, PPNet achieves an average improvement of 0.7% in AUC over multi-scenario recommendation methods in KuaiRand. However, their performance remains

inconsistent across datasets, particularly in those with severe distribution shifts, reflecting their limited capacity to handle highly heterogeneous scenarios. (5) DAR achieves consistent improvements over all baseline models, demonstrating its ability to balance cross-scenario alignment and scenario-specific learning. By incorporating a prototype layer and leveraging Optimal Transport, DAR minimizes the discrepancies between scenario distributions, facilitating efficient knowledge transfer. Additionally, the integration of scenario-specific expert networks allows DAR to retain scenario-specific features while benefiting from shared representations. On average, DAR outperforms the best-performing baseline models by 1.3% in AUC and 2.0% in LogLoss, with particularly strong results on imbalanced datasets like Douban.

Hence, our DAR well-addresses the challenges of balancing global and local learning in multi-scenario recommendation tasks. By combining prototype-based alignment with scenario-specific customization, DAR achieves consistent improvements across diverse datasets, establishing itself as a robust and state-of-the-art framework.

6.2 Ablation Study

In this section, we conduct ablation experiments to evaluate the contributions of different modules in DAR and analyze the effects of various design choices.

6.2.1 Different modules of DAR. Table 4 presents the ablation results to evaluate the contributions of different modules in DAR. Each row corresponds to a specific module being removed or replaced: (1) **w/o Re-representation**: This removes the re-representation process, directly using the basic representations for prediction. (2) **w/o $\mathcal{L}_{usage}$** : This removes the usage loss, which supervises the effective alignment between representations and prototypes. (3) **w/o $\mathcal{L}_{orth}$** : This removes the orthogonalization loss, which ensures the independence of global prototypes.

The results demonstrate that removing any module leads to a noticeable performance drop across all datasets. Specifically, removing the re-representation process results in the largest performance degradation, as it prevents the model from effectively disentangling shared and specific information. Removing the orthogonalization loss also significantly impacts performance, especially on the Douban and KuaiRand datasets, highlighting its importance in ensuring prototype independence.

6.2.2 Different prototypes construction strategies. Table 5 compares the impact of different strategies for constructing prototype spaces: (1) **Random**: Prototypes are initialized randomly without leveraging data structure. (2) **Sum-pooling**: Prototypes are computed as the sum of sample representations in the same cluster. (3) **K-Means**: Prototypes are constructed using K-Means clustering, which groups similar representations together.

The results indicate that K-Means outperforms the other two strategies across all datasets. While Sum-pooling performs better than Random, it is still inferior to K-Means due to its inability to model the inherent structure of the data. The significant improvement of K-Means highlights its effectiveness in creating meaningful and representative prototype spaces.

Table 2: AUC performance of different models. Higher values indicate better performance. "S1" indicates the results for scenario 1, and "ALL" represents results on the entire dataset. The best performance is shown in bold, and the second-best performance is shown with an underline. All results are the average of five runs, each with a different random seed.

Dataset	Scenario	ShareBottom	MMOE	PLE	STAR	SAR-Net	M2M	AdaSparse	CausalInt	EPNet	PPNet	M3oE	DAR
Douban	S1	0.6713	0.6892	0.6550	0.6827	0.6898	0.6782	0.6851	0.6602	0.6823	0.6865	<u>0.6902</u>	0.6983
	S2	0.6791	0.6824	0.6667	0.6878	0.6801	0.6872	0.6804	0.6853	0.6902	0.6842	<u>0.6931</u>	0.7004
	S3	0.8042	0.8101	0.7987	0.8112	0.8123	0.8089	0.8132	0.7998	0.8037	0.8148	<u>0.8021</u>	<u>0.8125</u>
	ALL	0.7794	0.7803	0.7765	0.7811	0.7809	0.7784	0.7815	0.7806	0.7823	<u>0.7829</u>	0.7817	0.7856
Ali-CCP	S1	0.5984	0.6063	0.6015	0.6001	0.6082	0.6047	0.6031	0.5962	0.6096	0.6038	0.6101	0.6108
	S2	0.5512	0.5584	0.5720	0.5659	0.5701	0.5627	0.5589	0.5635	<u>0.5791</u>	0.5744	0.5682	0.5827
	S3	0.6269	0.6302	0.6317	0.6373	0.6344	0.6249	0.6301	0.6287	0.6336	0.6264	0.6299	0.6356
	ALL	0.6187	0.6211	0.6209	<u>0.6213</u>	0.6110	0.6204	0.6184	0.6175	0.6157	0.6144	0.6207	0.6251
KuaiRand	S1	0.6891	0.6802	0.6954	0.6843	0.7010	0.6992	0.6759	0.6901	0.6973	<u>0.7025</u>	0.6835	0.7038
	S2	0.7111	0.7250	0.7039	0.7202	0.7147	0.7250	0.7215	0.7144	0.7233	0.7271	0.7152	<u>0.7264</u>
	S3	0.7742	0.7821	0.7785	0.7850	0.7719	0.7776	0.7734	0.7813	0.7792	0.7844	0.7838	0.7901
	S4	0.7252	0.7279	0.7193	<u>0.7305</u>	0.7234	0.7291	0.7288	0.7263	0.7315	0.7270	<u>0.7320</u>	0.7332
	S5	0.8389	0.8513	0.8477	0.8500	0.8406	0.8442	0.8497	0.8468	0.8542	0.8518	<u>0.8540</u>	0.8535
	ALL	0.7612	0.7654	0.7583	0.7680	0.7714	0.7658	0.7685	0.7593	0.7704	0.7718	<u>0.7731</u>	0.7785
Industry	S1	0.7012	0.7028	0.6954	0.7101	0.7083	0.7095	0.6967	0.7042	0.7070	<u>0.7112</u>	0.7008	0.7130
	S2	0.6652	0.6785	0.6604	0.6716	0.6739	0.6701	0.6770	0.6725	0.6795	<u>0.6724</u>	0.6846	0.6891
	S3	0.7502	0.7545	0.7480	0.7516	0.7537	<u>0.7556</u>	0.7522	0.7570	0.7499	0.7544	0.7513	0.7562
	ALL	0.7143	0.7189	0.7105	0.7204	0.7152	0.7143	0.7169	0.7137	0.7187	0.7196	<u>0.7209</u>	0.7240

Table 3: Logloss performance of different models. Lower values indicate better performance.

Dataset	Scenario	ShareBottom	MMOE	PLE	STAR	SAR-Net	M2M	AdaSparse	CausalInt	EPNet	PPNet	M3oE	DAR
Douban	S1	0.5671	0.5623	0.5704	0.5642	0.5620	0.5657	0.5631	0.5698	0.5645	0.5615	<u>0.5602</u>	0.5561
	S2	0.5098	0.5087	0.5104	0.5067	0.5095	0.5071	0.5092	0.5075	0.5061	0.5082	<u>0.5052</u>	0.5014
	S3	0.5432	0.5419	0.5462	0.5413	0.5407	0.5425	0.5400	0.5451	0.5437	<u>0.5385</u>	0.5443	0.5372
	ALL	0.5386	0.5381	0.5397	0.5364	0.5370	0.5392	0.5360	0.5375	0.5348	0.5342	<u>0.5340</u>	0.5337
Ali-CCP	S1	0.1808	0.1697	0.1793	0.1703	0.1771	0.1781	0.1789	0.1712	0.1665	0.1685	0.1697	0.1659
	S2	0.2179	0.2175	0.2047	0.2059	0.2051	0.2067	0.2070	0.2063	<u>0.2037</u>	0.2043	0.2055	0.1931
	S3	0.1717	0.1702	0.1711	<u>0.1690</u>	0.1694	0.1728	0.1706	0.1721	<u>0.1724</u>	0.1714	0.1698	0.1685
	ALL	0.1869	0.1783	0.1812	<u>0.1741</u>	0.1845	0.1881	0.1857	0.1873	0.1775	0.1761	0.1749	0.1736
KuaiRand	S1	0.3867	0.3821	0.3895	0.3846	0.3814	0.3857	0.3832	0.3880	0.3828	0.3810	<u>0.3802</u>	0.3782
	S2	0.6173	0.6148	0.6205	0.6157	0.6145	0.6169	0.6151	0.6192	0.6154	0.6139	<u>0.6134</u>	0.6119
	S3	0.5598	0.5581	0.5612	0.5570	0.5564	0.5589	0.5567	0.5603	0.5574	0.5556	0.5562	0.5541
	S4	0.6339	0.6312	0.6354	0.6303	0.6295	0.6328	0.6317	0.6342	0.6309	0.6291	<u>0.6285</u>	0.6277
	S5	0.3724	0.3685	0.3750	0.3697	0.3681	0.3709	0.3693	0.3732	0.3690	0.3678	<u>0.3664</u>	0.3650
	ALL	0.5663	0.5628	0.5679	0.5618	0.5612	0.5641	0.5625	0.5655	0.5623	0.5610	0.5613	0.5604
Industry	S1	0.1350	<u>0.1302</u>	0.1343	0.1336	0.1321	0.1338	0.1324	0.1349	0.1317	0.1314	0.1329	0.1286
	S2	0.1742	<u>0.1727</u>	0.1734	<u>0.1704</u>	0.1715	0.1725	0.1719	0.1737	0.1712	0.1709	0.1721	0.1693
	S3	0.1261	0.1245	0.1252	0.1233	<u>0.1213</u>	0.1240	0.1231	0.1257	0.1224	0.1218	0.1219	0.1203
	ALL	0.1648	0.1637	0.1642	0.1629	<u>0.1624</u>	0.1633	0.1626	0.1644	0.1619	0.1615	<u>0.1608</u>	0.1594

Table 4: Ablation results with AUC performance.

	Douban	Ali-CCP	KuaiRand	Industry
w/o Re-representation	0.7655	0.6032	0.7496	0.7013
w/o $\mathcal{L}_{usage}$	0.7693	0.6085	0.7609	0.7071
w/o $\mathcal{L}_{orth}$	0.7833	0.6239	0.7745	0.7212
DAR	0.7856	0.6251	0.7785	0.7240

Table 5: Different prototypes construction strategies.

	Douban	Ali-CCP	KuaiRand	Industry
Random	0.7802	0.6217	0.7728	0.7209
Sum-pooling	0.7834	0.6230	0.7749	0.7234
K-Means (ours)	0.7856	0.6251	0.7785	0.7240

Table 6: Different re-representation learning methods.

	Douban	Ali-CCP	KuaiRand	Industry
w/o $\mathcal{L}_{usage}$	0.7693	0.6085	0.7609	0.7071
L2	0.7581	0.6067	0.7614	0.7052
OT	0.7754	0.6189	0.7736	0.7193
UOT (ours)	0.7856	0.6251	0.7785	0.7240

6.2.3 Different re-representation learning methods. Table 6 evaluates the performance of DAR with different re-representation learning methods: (1) **w/o $\mathcal{L}_{usage}$** : The usage loss is removed, meaning no pseudo-labels guide the re-representation process. (2) **L2**: Re-representation learning is supervised using an L2 loss directly

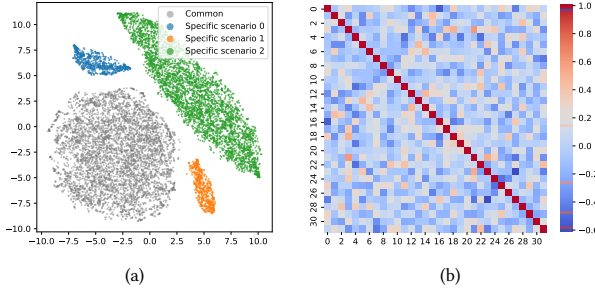


Figure 4: (a) The distribution of re-representations of Douban dataset was visualized by using t-SNE for dimensionality reduction. (b) The heatmap of the prototypes cosine similarity.

between representations and prototypes. (3) OT: Optimal Transport (OT) is used without balancing the distributional mismatch. (4) UOT (ours): Unbalanced Optimal Transport (UOT) is used to handle distributional imbalance.

The results show that UOT achieves the best performance across all datasets, demonstrating its ability to handle the inherent imbalance between individual representations and prototypes. While OT performs better than L2, it is less effective than UOT due to its inability to adapt to distributional differences. Removing the usage loss (w/o $\mathcal{L}_{usage}$) leads to noticeable performance degradation, further emphasizing the importance of supervision during the re-representation process.

6.3 Visualization

To provide an overview of the effectiveness of DAR in multi-scenario learning, we visualize the scenario-shared re-representation $\mathbf{H}$ and scenario-specific re-representation $\mathbf{R}$. Fig. 4 (a) shows the distribution of samples in the prototype space and the specific scenario spaces on Douban dataset. It can be seen that DAR effectively leverages the global prototypes $\mathbf{P}$ to learn common knowledge across samples from different scenarios, while the scenario specific prototypes retain the unique features of each scenario. Moreover, in Douban, the data distribution varies significantly between scenarios. Nevertheless, DAR effectively works in both cases, demonstrating its capability to achieve multi-scenario knowledge transfer and feature alignment through the combination of OT and prototypes to improve prediction performance. Additionally, we provide heatmaps of the prototype cosine similarity, as shown in Fig. 4 (b), illustrating the similarity distribution among the prototypes. The redder the color, the higher the similarity between prototypes; the bluer the color, the lower the similarity. It shows that almost all scenario-shared prototypes are dissimilar to each other, indicating that the prototypes, as an intermediary distribution, effectively absorb and store information from different scenarios.

6.4 Hyper-Parameter Analysis

We analyzed three key hyperparameters in the DAR model: the number of prototypes K , and the loss weights γ_1 (for UOT-based supervision) and γ_2 (for orthogonalization). Figure 5 shows its impact on performance. In Figure 5(a), increasing K improves the

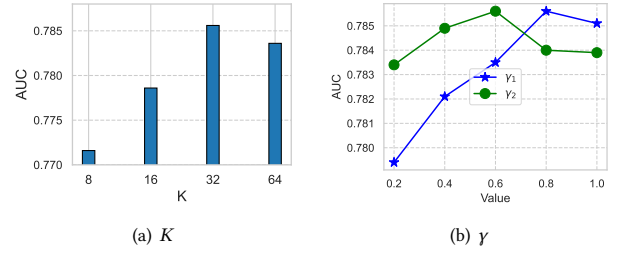


Figure 5: Hyper-parameter K and γ results of Douban.

Table 7: Online A/B results in three representative scenario in our APPs. We show the relative performance of DAR. The grey square brackets represent the 95% confidence intervals for online metrics.

Scenario	Watch Time	Like	Comment
Home page	+0.138% [0.01%, 0.29%]	+0.092% [−0.12%, 0.30%]	0.037% [−0.30%, 0.38%]
City page	+0.658% [0.06%, 1.26%]	+0.296% [−0.11%, 0.70%]	+0.587% [−0.34%, 1.49%]
Following page	+0.502% [0.05%, 1.05%]	+0.030% [−0.09%, 0.15%]	+0.322% [0.05%, 0.59%]

model’s ability to capture shared and scenario-specific features, but too many prototypes can increase complexity and introduce noise. Choosing an appropriate K is important for balancing expressiveness and efficiency. In Figure 5(b), γ_1 aligns representations with prototypes, while γ_2 encourages diversity among prototypes. Properly tuning both helps achieve effective alignment and prevents redundancy, boosting recommendation performance.

6.5 Online A/B Test

To further validate the effectiveness of DAR in practical applications, we conducted online A/B tests in three representative scenarios within the application: the Home page, the City page, and the Following page. The test results show that the DAR method significantly improved viewing time, likes, and comments in all scenarios. This indicates that DAR can effectively increase user viewing time and interaction frequency, further demonstrating its advantages and robustness in Multi-Scenario Recommendation tasks. These improvements not only confirm the effectiveness of DAR in cross-scenario knowledge transfer and feature alignment but also showcase its broad potential and adaptability in real-world applications.

7 Conclusion

In this paper, we proposed the Distribution-aware Re-representations (DAR) method to tackle the challenges of Multi-Scenario Recommendation (MSR), where varying user behaviors lead to heterogeneous yet interrelated data distributions. DAR leverages distribution-aware prototype spaces and disentangled re-representations to effectively capture both global and scenario-specific information. Through extensive offline experiments on multiple datasets and online A/B tests, DAR demonstrated significant improvements in

recommendation performance, showcasing its effectiveness and scalability for real-world applications.

References

- [1] Ceren Budak, Anitha Kannan, Rakesh Agrawal, and Jan Pedersen. 2014. Inferring user interests from microblogs. *AAAI ICWSM* (2014).
- [2] Rich Caruana. 1997. Multitask learning. *Machine learning* 28, 1 (1997), 41–75.
- [3] Jianxin Chang, Chenbin Zhang, Yiqun Hui, Dewei Leng, Yanan Niu, Yang Song, and Kun Gai. 2023. Pepnet: Parameter and embedding personalized network for infusing with personalized prior information. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 3795–3804.
- [4] Yuting Chen, Yanshi Wang, Yabo Ni, An-Xiang Zeng, and Lanfen Lin. 2020. Scenario-aware and Mutual-based approach for Multi-scenario Recommendation in E-Commerce. In *2020 International Conference on Data Mining Workshops (ICDMW)*. IEEE, 127–135.
- [5] Lenaic Chizat, Gabriel Peyré, Bernhard Schmitzer, and François-Xavier Vialard. 2018. Scaling algorithms for unbalanced optimal transport problems. *Math. Comp.* 87, 314 (2018), 2563–2609.
- [6] Paul Covington, Jay Adams, and Emre Sargin. 2016. Deep neural networks for youtube recommendations. In *Proceedings of the 10th ACM conference on recommender systems*. 191–198.
- [7] Chao Cui, Shisong Tang, Fan Li, Jiechao Gao, and Hechang Chen. 2025. Calibrating Video Watch-time Predictions with Credible Prototype Alignment. In *International Conference on Machine Learning*. PMLR, 11563–11584.
- [8] Marco Cuturi. 2013. Sinkhorn distances: Lightspeed computation of optimal transport. In *Advances in Neural Information Processing Systems*. 2292–2300.
- [9] Xin Dai, Xiangnan Kong, and Tian Guo. 2020. EPNet: Learning to exit with flexible multi-branch network. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*. 235–244.
- [10] Chenjiao Feng, Jiye Liang, Peng Song, and Zhiqiang Wang. 2020. A fusion collaborative filtering method for sparse data in recommender systems. *Information Sciences* 521 (2020), 365–379.
- [11] Ian Goodfellow. 2016. Deep learning.
- [12] James A Hanley and Barbara J McNeil. 1982. The meaning and use of the area under a receiver operating characteristic (ROC) curve. *Radiology* (1982).
- [13] Keli Hu, Wenping Chen, YuanZe Sun, Xiaozhao Hu, Qianwei Zhou, and Zirui Zheng. 2023. PPNet: Pyramid pooling based network for polyp segmentation. *Computers in Biology and Medicine* 160 (2023), 107028.
- [14] Diederik P Kingma and Jimmy Ba. 2014. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980* (2014).
- [15] Hyeyoung Ko, Suyeon Lee, Yoonseo Park, and Anna Choi. 2022. A survey of recommendation systems: recommendation models, techniques, and application fields. *Electronics* 11, 1 (2022), 141.
- [16] Fan Li, Jiazhen Huang, Shisong Tang, Bing Han, Huafeng Cao, Haochen Sui, Ting Xu, and Xiaoyu Kang. 2025. Contrastive Prototype Framework for Calibrating Video Recommendation. In *Proceedings of the 33rd ACM International Conference on Multimedia*. 6373–6382.
- [17] Fan Li, Xu Si, Shisong Tang, Dingmin Wang, Kunyan Han, Bing Han, Guorui Zhou, Yang Song, and Hechang Chen. 2024. Contextual distillation model for diversified recommendation. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 5307–5316.
- [18] Pengcheng Li, Runze Li, Qing Da, An-Xiang Zeng, and Lijun Zhang. 2020. Improving multi-scenario learning to rank in e-commerce by exploiting task relationships in the label space. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*. 2605–2612.
- [19] Xiaopeng Li, Jingtong Gao, Pengyue Jia, Yichao Wang, Wanyu Wang, Yejing Wang, Yuhao Wang, Huifeng Guo, and Ruiming Tang. 2024. Scenario-Wise Rec: A Multi-Scenario Recommendation Benchmark. *arXiv preprint arXiv:2412.17374* (2024).
- [20] Xiaopeng Li, Fan Yan, Xiangyu Zhao, Yichao Wang, Bo Chen, Huifeng Guo, and Ruiming Tang. 2023. Hamur: Hyper adapter for multi-domain recommendation. In *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management*. 1268–1277.
- [21] Xingyu Lu, Tianke Zhang, Chang Meng, Xiaobei Wang, Jinpeng Wang, Yi-Fan Zhang, Shisong Tang, Changyi Liu, Haojie Ding, Kaiyu Jiang, et al. 2025. Vlm as policy: Common-law content moderation framework for short video platform. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2*. 4682–4693.
- [22] Jiaqi Ma, Zhe Zhao, Xinyang Yi, Jilin Chen, Lichan Hong, and Ed H Chi. 2018. Modeling task relationships in multi-task learning with multi-gate mixture-of-experts. In *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining*. 1930–1939.
- [23] Wentao Ning, Xiao Yan, Weiwen Liu, Reynold Cheng, Rui Zhang, and Bo Tang. 2023. Multi-domain recommendation with embedding disentangling and domain alignment. In *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management*. 1917–1927.
- [24] Filippo Santambrogio. 2015. Optimal transport for applied mathematicians. *Birkhäuser*, NY 55, 58–63 (2015), 94.
- [25] Qijie Shen, Wanjie Tao, Jing Zhang, Hong Wen, Zulong Chen, and Quan Lu. 2021. SAR-Net: A scenario-aware ranking network for personalized fair recommendation in hundreds of travel scenarios. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*. 4094–4103.
- [26] Xiang-Rong Sheng, Liqin Zhao, Guorui Zhou, Xinyao Ding, Binding Dai, Qiang Luo, Siran Yang, Jingshan Lv, Chi Zhang, Hongbo Deng, et al. 2021. One model to serve all: Star topology adaptive recommender for multi-domain ctr prediction. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*. 4104–4113.
- [27] Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. 2014. Dropout: A simple way to prevent neural networks from overfitting. *The Journal of Machine Learning Research* 15, 1 (2014), 1929–1958.
- [28] Hongyan Tang, Junning Liu, Ming Zhao, and Xudong Gong. 2020. Progressive layered extraction (ple): A novel multi-task learning (mtl) model for personalized recommendations. In *Proceedings of the 14th ACM Conference on Recommender Systems*. 269–278.
- [29] Shisong Tang, Qing Li, Xiaoteng Ma, Ci Gao, Dingmin Wang, Yong Jiang, Qian Ma, Aoyang Zhang, and Hechang Chen. 2022. Knowledge-based temporal fusion network for interpretable online video popularity prediction. In *Proceedings of the ACM Web Conference 2022*. 2879–2887.
- [30] Shisong Tang, Qing Li, Dingmin Wang, Ci Gao, Wentao Xiao, Dan Zhao, Yong Jiang, Qian Ma, and Aoyang Zhang. 2023. Counterfactual video recommendation for duration debiasing. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 4894–4903.
- [31] Yu Tian, Bofang Li, Si Chen, Xubin Li, Hongbo Deng, Jian Xu, Bo Zheng, Qian Wang, and Chenliang Li. 2023. Multi-Scenario Ranking with Adaptive Feature Learning. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 517–526.
- [32] Yichao Wang, Huifeng Guo, Bo Chen, Weiwen Liu, Zhirong Liu, Qi Zhang, Zhicheng He, Hongkun Zheng, Weiwei Yao, Muyu Zhang, et al. 2022. Causalint: Causal inspired intervention for multi-scenario recommendation. In *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 4090–4099.
- [33] Yuhao Wang, Xiangyu Zhao, Bo Chen, Qidong Liu, Huifeng Guo, Huanshuo Liu, Yichao Wang, Rui Zhang, and Ruiming Tang. 2023. PLATE: A prompt-enhanced paradigm for multi-scenario recommendations. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 1498–1507.
- [34] Binrui Wu, Shisong Tang, Fan Li, Bing Han, Chang Meng, Jingyu Xiao, and Jiechao Gao. 2025. Aligning and Balancing ID and Multimodal Representations for Recommendation. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2*. 5029–5038.
- [35] Ziyuan Xia, Anchen Sun, Jingyi Xu, Yuanzhe Peng, Rui Ma, and Minghui Cheng. 2024. Contemporary Recommendation Systems on Big Data and Their Applications: A Survey. *IEEE Access* (2024).
- [36] Xu Xie, Fei Sun, Zhaoyang Liu, Shiwen Wu, Jinyang Gao, Jiandong Zhang, Bolin Ding, and Bin Cui. 2022. Contrastive learning for sequential recommendation. In *2022 IEEE 38th international conference on data engineering (ICDE)*. IEEE.
- [37] Yulin Xu, Chao Cui, Shisong Tang, Fan Li, Bing Han, Huafeng Cao, Jiechao Gao, and Hechang Chen. 2025. Leveraging Label Distributions as Anchors to Enhance Video Recommendation. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2*. 5129–5138.
- [38] Xuanhua Yang, Xiaoyu Peng, Penghui Wei, Shaoguo Liu, Liang Wang, and Bo Zheng. 2022. Adaspase: Learning adaptively sparse structures for multi-domain click-through rate prediction. In *Proceedings of the 31st ACM International Conference on Information & Knowledge Management*. 4635–4639.
- [39] Qianqian Zhang, Xinru Liao, Quan Liu, Jian Xu, and Bo Zheng. 2022. Leaving no one behind: A multi-scenario multi-task meta learning approach for advertiser modeling. In *Proceedings of the Fifteenth ACM International Conference on Web Search and Data Mining*. 1368–1376.
- [40] Zijian Zhang, Shuchang Liu, Jiaao Yu, Qingpeng Cai, Xiangyu Zhao, Chunxu Zhang, Ziru Liu, Qidong Liu, Hongwei Zhao, Lantao Hu, et al. 2024. M3oe: Multi-domain multi-task mixture-of experts recommendation framework. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 893–902.
- [41] Liwei Zhao, Dong Liu, Yanmin Xu, Zhiyong Cheng, and Haifeng Liu. 2020. A Graphical and Attentional Framework for Dual-Target Cross-Domain Recommendation. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*. ACM, 229–238.
- [42] Guorui Zhou, Na Mou, Ying Fan, Qi Pi, Weijie Bian, Chang Zhou, Xiaoqiang Zhu, and Kun Gai. 2019. Deep interest evolution network for click-through rate prediction. In *Proceedings of the AAAI conference on artificial intelligence*, Vol. 33. 5941–5948.
- [43] Jie Zhou, Xianshuai Cao, Wenhao Li, Lin Bo, Kun Zhang, Chuan Luo, and Qian Yu. 2023. Hinet: Novel multi-scenario & multi-task learning with hierarchical information extraction. In *2023 IEEE 39th International Conference on Data Engineering (ICDE)*. IEEE, 2969–2975.