

# Next Interest Flow: A Generative Pre-training Paradigm for Recommender Systems by Modeling All-domain Movelines

Chen Gao  
gaochen.gao@taobao.com  
Alibaba Group  
Hangzhou, China

Lv Shao  
shaolv.sl@taobao.com  
Alibaba Group  
Hangzhou, China

Zixin Zhao  
foriye.zzx@taobao.com  
Alibaba Group  
Hangzhou, China

Tong Liu  
yingmu@taobao.com  
Alibaba Group  
Hangzhou, China

## Abstract

Click-Through Rate (CTR) prediction, a cornerstone of modern recommender systems, has been dominated by discriminative models that react to past user behavior rather than proactively modeling user intent. Existing generative paradigms attempt to address this but suffer from critical limitations: Large Language Model (LLM) based methods create a semantic mismatch by forcing e-commerce signals into a linguistic space, while ID-based generation is constrained by item memorization and cold-start issues. To overcome these limitations, we propose a novel generative pre-training paradigm. Our model learns to predict the Next Interest Flow, a dense vector sequence representing a user’s future intent, while simultaneously modeling its internal Interest Diversity and Interest Evolution Velocity to ensure the representation is both rich and coherent. However, this two-stage approach introduces a critical objective mismatch between the generative and discriminative stages. We resolve this via a bidirectional alignment strategy, which harmonizes the two stages through cross-stage weight initialization and a dynamic Semantic Alignment Module for fine-tuning. Additionally, we enhance the underlying discriminative model with a Temporal Sequential Pairwise (TSP) mechanism to better capture temporal causality. We present the All-domain Moveline Evolution Network (AMEN), a unified framework implementing our entire pipeline. Extensive offline experiments validate AMEN’s superiority over strong baselines, and a large-scale online A/B test demonstrates its significant real-world impact, delivering substantial improvements in key business metrics.

## 1 Introduction

Click-Through Rate (CTR) prediction, a cornerstone of modern recommender systems, has long been dominated by discriminative models, such as DIN [17], DIEN [16], DSIN [5], SIM [9], DIAN [12] and CCN [6]. These models learn from sequences of user-item interactions to predict the likelihood of a click. However, their fundamental limitation lies in their reactive nature: they are trained to explain past user feedback rather than to proactively model a user’s future intent. This approach often ignores the rich contextual signals present in the all-domain user moveline, which is the user’s complete journey across heterogeneous scenes like browsing, searching, and engaging with promotions.

The need for proactive, forward-looking prediction motivates a shift towards generative paradigms. However, existing approaches

fall short. The first paradigm, leveraging Large Language Models (LLMs), such as P5 [7], M6-Rec [3], HLLM [1], translates user interactions into text to predict a future item’s description. This creates a semantic mismatch, forcing nuanced e-commerce signals into a linguistic space, and requires an inefficient post-generation retrieval step to map text back to an actual item. The second paradigm, which directly generates item IDs, such as TIGER [10], HSTU [14], OneRec [4, 15], frames the task as item memorization rather than semantic understanding. This approach relies on vector quantization methods [11, 13] to create a compressed codebook to handle massive ID vocabularies, and struggles with long-tail items and cold-start challenges.

To overcome these limitations, we propose a novel generative pre-training paradigm centered on a new concept: the Next Interest Flow. Instead of generating text or a discrete ID, our model learns to predict a dense vector sequence that directly represents the user’s evolving future interests within the e-commerce semantic space. This flow is not used for direct item retrieval but serves as a powerful, predictive feature for the downstream discriminative CTR model. This design sidesteps the shortcomings of prior work: it avoids the semantic mismatch of LLM methods and eliminates the vocabulary constraints and cold-start issues inherent to ID-based generation. Furthermore, we model the flow’s internal Interest Diversity and Interest Evolution Velocity, providing a diverse and coherent forecast of user intent.

However, this two-stage process introduces a critical challenge: an objective mismatch between the generative pre-training (predicting the next interest vector) and the discriminative fine-tuning (predicting the click-through rate on a specific target item). We resolve this through a bidirectional alignment strategy. First, we initialize the generator with the fundamental embedding weights of the basic discriminative model, ensuring it starts from a semantically relevant parameter space. Second, we introduce a Semantic Alignment Module during fine-tuning to dynamically reconcile the predicted flow with the specific context of the target item, ensuring the pre-trained knowledge is precisely adapted for the final click prediction.

Finally, we recognize that the power of our generative model is built upon a robust discriminative foundation. To this end, we enhance the underlying CTR model with a Temporal Sequential Pairwise (TSP) learning mechanism. This self-supervised task is designed to instill an awareness of temporal causality, enabling the

model to better understand how prior actions within the moveline causally influence subsequent user feedback.

We present the All-domain Moveline Evolution Network (AMEN), a unified framework that implements our entire pipeline. AMEN integrates our novel generative paradigm, the bidirectional alignment strategy, and the enhanced TSP-based discriminative model. It first pre-trains a generator to produce the Next Interest Flow, which is then consumed by the final CTR model for prediction. The contributions of this work are as follows:

- We propose a new generative pre-training paradigm for recommendation that predicts a user's Next Interest Flow, a dense vector representation of future intent, overcoming the core limitations of existing LLM-based and ID-based generative models.
- We identify and solve the objective mismatch inherent in the two-stage approach via a bidirectional alignment strategy, consisting of cross-stage weight initialization and a dynamic Semantic Alignment Module.
- We introduce the Temporal Sequential Pairwise (TSP) mechanism to enhance the underlying discriminative model, enabling it to better capture temporal causality within the user's all-domain moveline.
- We present AMEN, a unified framework implementing our proposed paradigm. Extensive offline experiments validate its superiority over strong baselines, and a large-scale online A/B test demonstrates its significant real-world business impact.

## 2 Proposed Method: The AMEN Framework

We propose the All-domain Moveline Evolution Network (AMEN), a two-stage framework for CTR prediction. The framework is designed to first learn a generative model of user intent from their all-domain moveline, and then use this knowledge to enhance a discriminative prediction model. The two stages are as follows:

- **Stage 1: Generative Pre-training.** A Transformer-based decoder,  $G_\phi$ , is pre-trained to predict the Next Interest Flow, which is a dense vector representing a user's future interest, based on their historical moveline.
- **Stage 2: Discriminative Fine-tuning.** A discriminative CTR model,  $F_\theta$ , is trained for the final prediction task. It utilizes the frozen, pre-trained generator  $G_\phi$  to produce the Next Interest Flow as a key input feature. The model architecture is designed to explicitly handle this generative input.

The overall architecture is illustrated in Figure 1.

### 2.1 Stage 1: Generative Pre-training

**Next Interest Flow.** We define the Next Interest Flow,  $\mathbf{F} \in \mathbb{R}^d$ , as a dense vector representation of a user's evolving interest, generated from their moveline  $\vec{\mathcal{M}}$ .

**Model and Objective.** The generative model  $G_\phi$  is a Transformer-based decoder, trained to autoregressively predict the future Next Interest Flow. For a given historical moveline ending at a specific time  $t_0$ , the model sequentially generates  $T$  future flow states, denoted as  $\hat{\mathbf{F}}_{\geq t_0} = (\hat{\mathbf{f}}_{t_0}, \dots, \hat{\mathbf{f}}_{t_0+T-1})$ .

The training employs a teacher forcing strategy. At each future timestep  $t$  within the prediction window  $[t_0, t_0 + T - 1]$ , the model  $G_\phi$  receives the initial moveline  $\vec{\mathcal{M}}_{< t_0}$  to predict the next flow state  $\hat{\mathbf{f}}_t$ . The objective is to maximize the similarity, implemented as the dot product, between each predicted flow state  $\hat{\mathbf{f}}_t$  and its corresponding ground-truth item embedding  $\mathbf{e}_t$  (the positive sample), while minimizing it against negative samples.

The overall pre-training objective minimizes the sum of InfoNCE [8] losses across the entire prediction window:

$$\mathcal{L}_{G_\phi} = - \sum_{t=t_0}^{t_0+T-1} \log \frac{\exp(\text{sim}(\hat{\mathbf{f}}_t, \mathbf{e}_t))}{\sum_{\mathbf{e}_i \in \{\mathbf{e}_t\} \cup \mathcal{N}_t} \exp(\text{sim}(\hat{\mathbf{f}}_t, \mathbf{e}_i))} \quad (1)$$

where  $\hat{\mathbf{f}}_t = G_\phi(\vec{\mathcal{M}}_{< t_0})$ ,  $\mathcal{N}_t$  is the set of negative samples for timestep  $t$ ,  $\text{sim}(\mathbf{u}, \mathbf{v}) = (\mathbf{u} \cdot \mathbf{v}) / \tau$  represents the dot product with a scaling temperature coefficient  $\tau$ .

**Interest Diversity.** The Next Interest Flow state vector  $\hat{\mathbf{f}}_t$  is not a monolithic point in semantic space, but is composed of outputs from multi attention heads within our Transformer-based generator. This architectural property provides a natural way to model a user's interest diversity.

Let our generator have  $H$  attention heads. The output of each head  $h$  is a sub-vector  $\mathbf{h} \in \mathbb{R}^{d_{head}}$ , which is designed to capture a different facet or subspace of the user's interest. The final flow state vector is the concatenation of these sub-vector outputs:

$$\hat{\mathbf{f}}_t = \text{Concat}(\mathbf{h}_1, \mathbf{h}_2, \dots, \mathbf{h}_H) \quad (2)$$

The total dimension of the flow state vector is thus  $d = H \times d_{head}$ .

To ensure that different heads learn distinct and non-redundant representations, we apply a diversity loss that encourages dissimilarity between the outputs of these heads. It is formulated as a repulsion loss based on the pairwise similarity between the head-specific sub-vectors:

$$\mathcal{L}_{div} = \sum_{t=t_0}^{t_0+T-1} \left( \frac{1}{\binom{H}{2}} \sum_{i=1}^{H-1} \sum_{j=i+1}^H (\text{sim}(\mathbf{h}_i, \mathbf{h}_j))^2 \right) \quad (3)$$

where the sum is over all pairs of distinct attention heads  $(i, j)$ . This loss penalizes heads that produce similar outputs, forcing the model to utilize its full representational capacity.

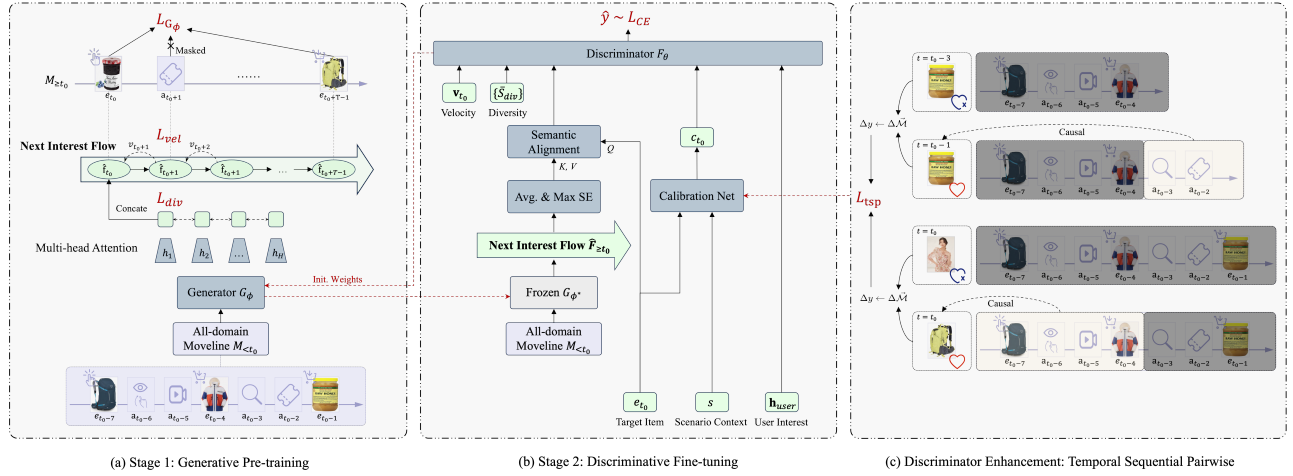
During inference, we compute a per-user diversity score,  $S_{div}(t) = 1 - \mathcal{L}_{div}(\hat{\mathbf{f}}_t)$ , which is used as a feature to distinguish between users seeking variety versus those seeking depth.

**Interest Evolution Velocity.** Beyond predicting the static interest at a future point, our framework can model the dynamics of interest evolution. We define the Interest Evolution Velocity at time  $t$  as the difference between continuous predicted flow states. This velocity vector captures both the direction and magnitude of the user's interest shift.

$$\mathbf{v}_t = \hat{\mathbf{f}}_t - \hat{\mathbf{f}}_{t-1} \quad (4)$$

We introduce an additional regularization term that encourages smoothness or consistency in the evolution, penalizing abrupt, random-like jumps. The velocity loss can be formulated as:

$$\mathcal{L}_{vel} = \sum_{t=t_0+1}^{t_0+T-1} \left\| (\hat{\mathbf{f}}_t - \hat{\mathbf{f}}_{t-1}) - (\hat{\mathbf{f}}_{t-1} - \hat{\mathbf{f}}_{t-2}) \right\|_2^2 \quad (5)$$



**Figure 1: The overall architecture of the All-domain Moveline Evolution Network (AMEN). (a) Stage 1: Generative Pre-training. (b) Stage 2: Discriminative Fine-tuning. (c) Discriminator Enhancement: Temporal Sequential Pairwise (TSP) Task.**

The velocity vector  $\mathbf{v}_t$  also serves as feature for downstream tasks, indicating periods of high user exploration.

The final loss for pre-training is a weighted sum of these objectives:

$$\mathcal{L}_{\text{stage1}} = \mathcal{L}_{G_\phi} + \alpha \cdot \mathcal{L}_{\text{div}} + \beta \cdot \mathcal{L}_{\text{vel}} \quad (6)$$

where  $\alpha, \beta$  are hyperparameters balancing multi objectives.

**Weight Initialization.** To align the parameter spaces of the two stages, the generative decoder  $G_\phi$  is initialized with the weights of a pre-trained discriminative base model before the contrastive learning begins.

## 2.2 Stage 2: Discriminative Fine-tuning

In the fine-tuning stage, the pre-trained generator  $G_\phi$  acts as a feature extractor, providing a forward-looking user's Next Interest Flow to the discriminative CTR model  $F_\theta$ . The key challenge is to effectively apply the generated flow for predicting the click-through rate on a single target item  $e_{t_0}$  at a given inference timestamp  $t_0$ .

**Input Features.** At inference time  $t_0$ , for a given moveline  $\vec{M}_{<t_0}$ , the frozen generator  $G_\phi$  produces three types of forward-looking features:

- (1) **Next Interest Flow:** The sequence of flow state vectors  $\hat{\mathbf{F}}_{\geq t_0} = (\hat{\mathbf{f}}_{t_0}, \dots, \hat{\mathbf{f}}_{t_0+T-1})$ .
- (2) **Interest Diversity:** Per-state diversity scores  $\{S_{\text{div}}(\hat{\mathbf{f}}_t)\}_{t=t_0}^{t_0+T-1}$  derived from the internal structure of each flow state.
- (3) **Interest Evolution Velocity:** The sequence of velocity vectors  $\{\mathbf{v}_t\}_{t=t_0+1}^{t_0+T-1}$ .

**Semantic Alignment Module.** The raw flow  $\hat{\mathbf{F}}_{\geq t_0}$  represents a general future trajectory. To reconcile this with the specific target item  $e_{t_0}$ , we employ a Semantic Alignment Module. The target item's embedding  $e_{t_0}$  acts as a query to attend to the sequence of flow states, producing a single, context-aware representation  $\mathbf{a}_{\text{flow}}$  that summarizes the most relevant future interests for this specific item:

$$\mathbf{a}_{\text{flow}} = \text{Attention}(\text{Query} = e_{t_0}, \text{Key} = \hat{\mathbf{F}}_{\geq t_0}, \text{Value} = \hat{\mathbf{F}}_{\geq t_0}) \quad (7)$$

This allows the model to dynamically decide whether the immediate future flow state ( $\hat{\mathbf{f}}_{t_0}$ ) or a more distant one is more indicative of a click.

**Final Prediction.** The aligned flow representation  $\mathbf{a}_{\text{flow}}$  is concatenated with other model features. These include:

- Traditional user interest representations  $\mathbf{h}_{\text{user}}$  (e.g., from MHTA [17] on historical behavior sequences).
- Aggregated generative features: We pool the diversity and velocity features, for instance, by taking the mean of the diversity scores  $\{\hat{S}_{\text{div}}\}$  and using the first velocity vector  $\mathbf{v}_{t_0}$ .
- Standard features like user profile and the target item embedding  $e_{t_0}$ .

This comprehensive combined vector is fed into a final merge MLP to produce a logit  $\hat{y}_{\text{main}}$ . The final prediction incorporates a calibration score  $c_{t_{\text{sp}}}$  from the auxiliary TSP task:

$$\hat{y} = \sigma(\hat{y}_{\text{main}} + c_{t_{\text{sp}}}) \quad (8)$$

## 2.3 Temporal Sequential Pairwise (TSP) Auxiliary Task

During fine-tuning, we incorporate the TSP task to explicitly model temporal causality within the moveline.

**Formulation.** The TSP task operates on pairs of samples. For a target sample  $(e_{t_0}, \vec{M}_{<t_0})$ , a diff sample  $(e_{t_1}, \vec{M}_{<t_1})$  with an opposite click label is selected from a different time point within window  $T$ . A Calibration Net within  $F_\theta$  calculates a moveline-based calibration score for both:  $c_{t_0}$  for the target and  $c_{t_1}$  for the diff sample. The TSP loss, a variant of BPR, maximizes the margin between these scores based on their click labels:

$$\mathcal{L}_{\text{tsp}} = -\frac{1}{|\mathcal{D}_{\text{paired}}|} \sum_{\mathcal{D}_{\text{paired}}} \log \sigma(\mathbb{I}(y_{t_1})(c_{t_1} - c_{t_0})) \quad (9)$$

where  $\mathbb{I}(y_{t_1})$  is an indicator function (+1 for a positive label, -1 for a negative label). The scalar  $c_{t_0}$  used in the inference is the calibration score computed for the target item.

The final loss for the fine-tuning stage is a weighted sum of the standard cross-entropy loss and the TSP loss:

$$\mathcal{L}_{\text{stage2}} = \mathcal{L}_{\text{CE}} + \lambda \cdot \mathcal{L}_{\text{TSP}} \quad (10)$$

where  $\lambda$  is a balancing hyperparameter.

The complete training procedure are detailed in Algorithm 1.

---

**Algorithm 1: AMEN Training Procedure**


---

**Input:** Dataset  $\mathcal{D}$ , base model weights  $\Theta_{\text{base}}$ , hyperparameters  $\alpha, \beta, \lambda, T, H, \tau$ .

**Output:**

- 1 Trained AMEN parameters  $\Theta_{\text{AMEN}} = (\phi^*, \theta^*)$ .

```

/* Stage 1: Generative Pre-training */
2 Initialize generator  $G_\phi$  with weights from  $\Theta_{\text{base}}$ ;
3 for each batch  $B \subset \mathcal{D}$  do
4   for each sample with moveline  $\vec{M}_{<t_0}$  and future items
      $\{e_t\}_{t=t_0}^{t_0+T-1}$  do
5     for  $t \in [t_0, t_0 + T - 1]$  do
6        $\hat{\mathbf{f}}_t \leftarrow G_\phi(\vec{M}_{<t_0})$ ;
7       Compute  $\mathcal{L}_{G_\phi}$  for  $\hat{\mathbf{f}}_t$  using Eq. (1);
8       Compute  $\mathcal{L}_{\text{div}}$  using  $\{\mathbf{h}\}$  of  $\hat{\mathbf{f}}_t$ ;
9       Compute  $\mathcal{L}_{\text{vel}}$  using  $\{\mathbf{v}_t\}$ ;
10    end
11     $\mathcal{L}_{\text{stage1}} \leftarrow \sum_t (\mathcal{L}_{G_\phi} + \alpha \mathcal{L}_{\text{div}} + \beta \mathcal{L}_{\text{vel}})$ ;
12    Update  $\phi$  using gradient descent on  $\mathcal{L}_{\text{stage1}}$ ;
13  end
14 end
15 Let the trained generator be  $G_{\phi^*}$ ;

/* Stage 2: Discriminative Fine-tuning */
16 Initialize discriminative model  $F_\theta$ ;
17 for each batch  $B \subset \mathcal{D}$  do
18   for each target sample  $(e_{t_0}, \vec{M}_{<t_0}, y_{t_0})$  do
19      $\hat{\mathbf{F}}_{\geq t_0} \leftarrow G_{\phi^*}(\vec{M}_{<t_0})$ ;
20      $\mathbf{a}_{\text{flow}} \leftarrow \text{Attention}(e_{t_0}, \hat{\mathbf{F}}_{\geq t_0}, \hat{\mathbf{F}}_{\geq t_0})$ ;
21      $\tilde{\mathbf{S}}_{\text{div}} \leftarrow \text{Mean}(\{\mathbf{S}_{\text{div}}(\hat{\mathbf{f}}_t)\})$ ;  $\mathbf{v}_{\text{first}} \leftarrow \mathbf{v}_{t_0}$ ;
22      $\mathbf{z} \leftarrow \text{Concat}(\mathbf{a}_{\text{flow}}, \mathbf{h}_{\text{user}}, \tilde{\mathbf{S}}_{\text{div}}, \mathbf{v}_{\text{first}}, e_{t_0}, \dots)$ ;
23      $\hat{y}_{\text{main}} \leftarrow F_\theta[\text{MLP}](\mathbf{z})$ ;
24      $c_{t_0} \leftarrow F_\theta[\text{CalibrationNet}](e_{t_0}, \vec{M}_{<t_0})$ ;
25     Sample diff pair  $(e_{t_1}, \vec{M}_{<t_1}, y_{t_1})$  and compute its
       score  $c_{t_1}$ ;
26      $\mathcal{L}_{\text{TSP}} \leftarrow -\log \sigma(\mathbb{I}(y_{t_1})(c_{t_1} - c_{t_0}))$ ;
27      $\hat{y} \leftarrow \sigma(\hat{y}_{\text{main}} + c_{t_0})$ ;
28      $\mathcal{L}_{\text{CE}} \leftarrow \text{CrossEntropy}(y_{t_0}, \hat{y})$ ;
29      $\mathcal{L}_{\text{stage2}} \leftarrow \mathcal{L}_{\text{CE}} + \lambda \mathcal{L}_{\text{TSP}}$ ;
30     Update  $\theta$  using gradient descent on  $\mathcal{L}_{\text{stage2}}$ ;
31   end
32 end
33 return  $(\phi^*, \theta^*)$ ;

```

---

### 3 Experiments

We conduct experiments to answer three key questions:

- RQ1: How does AMEN compare to state-of-the-art models?
- RQ2: What is the impact of each component in our framework?
- RQ3: Does AMEN bring improvements in the real-world online industry setting?

#### 3.1 Experimental Setup

**Dataset.** We use a large-scale industrial dataset from the Taobao user logs, as details in Table 1.

**Table 1: Statistics of the industrial dataset**

| Split    | Users  | Items | Instances |
|----------|--------|-------|-----------|
| Training | 180.7M | 21.9M | 6.7B      |
| Test     | 23.2M  | 7.7M  | 0.6B      |

**Baselines.** We compare AMEN against: 1) Discriminative models: Wide&Deep [2], DIN [17]; 2) Generative models: P5 [7] (LLM-based) and TIGER [10] (ID-based); 3) MEDN, a strong internal discriminative baseline.

**Metrics.** We use Area Under Curve (AUC) [17] for offline evaluation.

#### 3.2 Offline Results and Ablation Study (RQ1, RQ2)

Table 2 presents the main performance comparison and the ablation study. We start from the MEDN baseline and incrementally add our proposed components to show their effects.

**Table 2: Offline results and ablation study on the industrial dataset.  $\Delta\text{AUC}$  is relative to the MEDN baseline. The ablation study removes components from the full AMEN model.**

| Category       | Model                        | AUC ( $\Delta\text{AUC}$ ) |
|----------------|------------------------------|----------------------------|
| Baselines      | Wide&Deep                    | 0.7216 ( -0.05pt )         |
|                | DIN                          | 0.7410 ( -2.11pt )         |
|                | P5 (Generative)              | 0.7565 ( -0.56pt )         |
|                | MEDN (Baseline)              | 0.7621 ( - )               |
|                | TIGER (Generative)           | 0.7638 ( +0.17pt )         |
| Ours           | <b>AMEN (Full Model)</b>     | <b>0.7708 ( +0.87pt )</b>  |
|                | w/o Next Interest Flow (NIF) | 0.7694 ( +0.73pt )         |
|                | w/o TSP Mechanism            | 0.7683 ( +0.62pt )         |
|                | w/o Sem. Align.              | 0.7702 ( +0.81pt )         |
|                | w/o Weight Init.             | 0.7698 ( +0.77pt )         |
|                | w/o Diversity Loss           | 0.7697 ( +0.76pt )         |
| Ablation Study | w/o Velocity Loss            | 0.7699 ( +0.78pt )         |

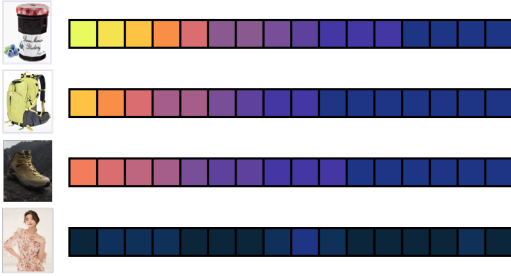
**Performance Comparison (RQ1).** AMEN significantly outperforms all baselines, achieving an AUC of 0.7708. This represents a substantial +0.87pt gain over the strong MEDN production baseline. Notably, AMEN outperforms generative paradigms: TIGER (ID-based) by +0.70pt and P5 (LLM-based) by +1.43pt. This confirms

the superiority of our Next Interest Flow approach, which avoids the pitfalls of item memorization and semantic mismatch.

**Ablation Study (RQ2).** The ablation study validates the effectiveness of each component. Removing the Next Interest Flow (NIF) features or the TSP Mechanism causes the two largest performance drops (0.14pt and 0.25pt respectively), highlighting them as the primary performance drivers. The bidirectional alignment strategy is also crucial: removing Semantic Alignment or Weight Initialization leads to AUC degradation. Finally, the auxiliary regularizers for Diversity and Velocity also provide positive contributions, confirming their role in learning a richer flow representation.

**Discussion on Component Contribution.** An interesting observation from the ablation is that removing the TSP mechanism causes a larger drop in AUC than removing the NIF features. This does not diminish the value of our generative paradigm but rather highlights two points: (1) The TSP task is a powerful, standalone enhancement for modeling temporal dynamics in complex user movelines. (2) The AMEN w/o NIF model still benefits from the TSP task, which explains why it remains a strong model. The greatest performance is achieved when both powerful components, the temporal-causal understanding from TSP and the proactive future prediction from NIF, are combined in the full AMEN model. Our online results further confirm this synergy.

**Analysis of the Next Interest Flow.** To intuitively demonstrate the richness and multiplexed nature of the NIF, we conduct a qualitative analysis. For the same user and historical moveline  $\mathcal{M}_{<t_0}$  shown in Figure 1, we generate the corresponding NIF  $\hat{\mathbf{F}}_{\geq t_0}$ . We then use the Semantic Alignment Module as a probe, feeding it different target items to see which parts of the NIF are activated. The resulting attention weights are visualized as a heatmap in Figure 2.

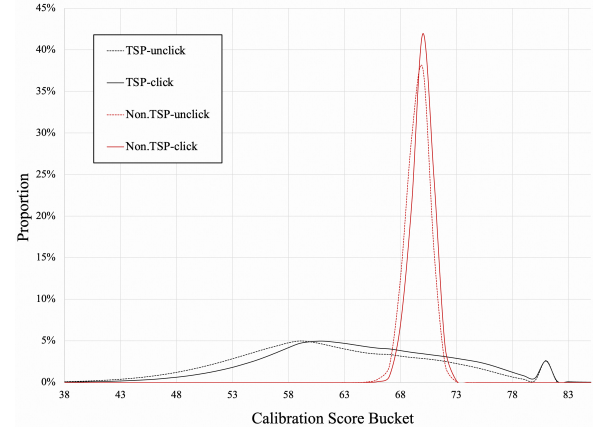


**Figure 2: Visualization of the information decoded from the Next Interest Flow.**

The result reveals that the NIF is not a monolithic prediction but a structured representation containing multiple, distinct interest signals. When probed with target items from the relative semantic category (the backpack and hiking boot), a consistent set of latent vectors within the NIF is activated, indicating that a specific subspace within the NIF has learned to represent an "outdoor activity" interest. When probed with a target item from a different category (the jam), a different primary set of vectors is activated. However, a slight overlap with the "outdoor" pattern is observable. This observation aligns with our *Interest Diversity* mechanism, where individual flow states, composed of multiple heads, can capture diverse interest and represent more general concepts (e.g., "trip supplies")

that are partially relevant to multiple categories. Finally, the diffuse attention for the dress confirms the NIF's context-specificity, as it generates strong signals only for interests consistent with the user's moveline.

**Analysis of the TSP Mechanism.** To better understand how the TSP mechanism enhances model performance, we visualize the distribution of its output: the calibration score  $c_{\text{tsp}}$ . We plot the probability density of these scores on the test set for two model variants: AMEN with TSP, and AMEN w/o TSP, both w/o NIF. For each variant, we analyze the score distributions for both positive (clicked) and negative (unclicked) samples.



**Figure 3: Probability density distributions of the TSP calibration score ( $c_{\text{tsp}}$ ).**

As illustrated in Figure 3, we draw two key conclusions:

(i) **Baseline Discrimination:** Both models learn to assign higher scores to positive samples, confirming a foundational discriminative ability.

(ii) **TSP Enhancement:** The TSP task significantly amplifies this effect. For AMEN with TSP, the distribution range is much wider and sparser, directly visualizing how TSP provides a stronger, more discriminative signal.

### 3.3 Online A/B Testing (RQ3)

We deploy AMEN in online A/B test against the MEDN industrial baseline on Taobao. As shown in Table 3, AMEN delivered improvements in key business metrics. Notably, it achieved a +11.6% lift in post-view CTCVR (a metric combining click and conversion) in the main Feeds scenario, confirming its substantial real-world value.

## 4 Conclusion

This paper introduced AMEN, a framework that advances CTR prediction by shifting from a reactive, discriminative paradigm to a proactive, generative one. Our core contributions are twofold. First, we propose generating a Next Interest Flow directly in the e-commerce semantic space, which overcomes the limitations of prior LLM-based and ID-based methods. Second, we identify and solve the objective mismatch problem in this two-stage process via a bidirectional alignment strategy: cross-stage weight initialization and semantic alignment. Our proposed TSP mechanism also proves

**Table 3: Online A/B testing results**

| Comparison                                 | CTCVR  | CTR    | GMV     |
|--------------------------------------------|--------|--------|---------|
| <i>Part 1: Discriminative Enhancements</i> |        |        |         |
| Feeds: AMEN w/o NIF vs. MEDN               | +11.6% | Sta.   | Sta.    |
| Floors: AMEN w/o NIF vs. MEDN              | +4.2%  | +20.6% | +20.1%  |
| <i>Part 2: Generative Paradigm</i>         |        |        |         |
| AMEN (Full) vs. w/o NIF                    | +2.28% | +0.98% | +11.24% |

CTCVR: Post-view Click-Through-Conversion Rate. Stable (Sta.) indicates a non-significant change.

to be a powerful tool for capturing temporal causality. Validated by extensive offline experiments and a large-scale online A/B test, AMEN establishes an effective new paradigm for proactive user interest modeling in recommender systems.

## References

- [1] Junyi Chen, Lu Chi, Bingyue Peng, and Zehuan Yuan. 2024. HLLM: Enhancing Sequential Recommendations via Hierarchical Large Language Models for Item and User Modeling. *arXiv:2409.12740 [cs.IR]* <https://arxiv.org/abs/2409.12740>
- [2] Heng-Tze Cheng, Levent Koc, Jeremiah Harmsen, Tal Shaked, Tushar Chandra, Hrishi Aradhye, Glen Anderson, Greg Corrado, Wei Chai, Mustafa Ispir, et al. 2016. Wide & deep learning for recommender systems. In *Proceedings of the 1st workshop on deep learning for recommender systems*. 7–10.
- [3] Zeyu Cui, Jianxin Ma, Chang Zhou, Jingren Zhou, and Hongxia Yang. 2022. M6-rec: Generative pretrained language models are open-ended recommender systems. *arXiv preprint arXiv:2205.08084* (2022).
- [4] Jiaxin Deng, Shiyao Wang, Kuo Cai, Lejian Ren, Qigen Hu, Weifeng Ding, Qiang Luo, and Guorui Zhou. 2025. OneRec: Unifying Retrieve and Rank with Generative Recommender and Iterative Preference Alignment. *arXiv:2502.18965 [cs.IR]* <https://arxiv.org/abs/2502.18965>
- [5] Yufei Feng, Fuyu Lv, Weichen Shen, Menghan Wang, Fei Sun, Yu Zhu, and Keping Yang. 2019. Deep session interest network for click-through rate prediction. *arXiv preprint arXiv:1905.06482* (2019).
- [6] Chen Gao, Zixin Zhao, Sihao Hu, Lv Shao, and Tong Liu. 2024. Collaborative Contrastive Network for Click-Through Rate Prediction. *arXiv:2411.11508 [cs.IR]* <https://arxiv.org/abs/2411.11508>
- [7] Shijie Geng, Shuchang Liu, Zuohui Fu, Yingqiang Ge, and Yongfeng Zhang. 2022. Recommendation as language processing (rlp): A unified pretrain, personalized prompt & predict paradigm (p5). In *Proceedings of the 16th ACM conference on recommender systems*. 299–315.
- [8] Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748* (2018).
- [9] Qi Pi, Xiaoqiang Zhu, Guorui Zhou, Yujing Zhang, Zhe Wang, Lejian Ren, Ying Fan, and Kun Gai. 2020. Search-based User Interest Modeling with Lifelong Sequential Behavior Data for Click-Through Rate Prediction. *CoRR abs/2006.05639* (2020). *arXiv:2006.05639*
- [10] Shashank Rajput, Nikhil Mehta, Anima Singh, Raghunandan Hulikal Keshavan, Trung Vu, Lukasz Heldt, Lichan Hong, Yi Tay, Vinh Tran, Jonah Samost, et al. 2023. Recommender systems with generative retrieval. *Advances in Neural Information Processing Systems* 36 (2023), 10299–10315.
- [11] Aaron van den Oord, Oriol Vinyals, and Koray Kavukcuoglu. 2017. Neural discrete representation learning. In *Proceedings of the 31st International Conference on Neural Information Processing Systems* (Long Beach, California, USA) (NIPS’17). Curran Associates Inc., Red Hook, NY, USA, 6309–6318.
- [12] Yaxian Xia, Yi Cao, Sihao Hu, Tong Liu, and Lingling Lu. 2023. Deep Intention-Aware Network for Click-Through Rate Prediction. In *Companion Proceedings of the ACM Web Conference 2023*. 533–537.
- [13] Jun Yin, Zhengxin Zeng, Mingzheng Li, Hao Yan, Chaozhuo Li, Weihao Han, Jianjin Zhang, Ruochen Liu, Hao Sun, Weiwei Deng, Feng Sun, Qi Zhang, Shirui Pan, and Senzhang Wang. 2025. Unleash LLMs Potential for Sequential Recommendation by Coordinating Dual Dynamic Index Mechanism. In *Proceedings of the ACM on Web Conference 2025* (Sydney NSW, Australia) (WWW ’25). Association for Computing Machinery, New York, NY, USA, 216–227. <https://doi.org/10.1145/3696410.3714866>
- [14] Jiaqi Zhai, Lucy Liao, Xing Liu, Yueming Wang, Rui Li, Xuan Cao, Leon Gao, Zhaojie Gong, Fangda Gu, Michael He, Yinghai Lu, and Yu Shi. 2024. Actions Speak Louder than Words: Trillion-Parameter Sequential Transducers for Generative Recommendations. *arXiv:2402.17152 [cs.LG]* <https://arxiv.org/abs/2402.17152>
- [15] Guorui Zhou, Jiaxin Deng, Jinghao Zhang, Kuo Cai, Lejian Ren, Qiang Luo, Qianqian Wang, Qigen Hu, Rui Huang, Shiyao Wang, Weifeng Ding, Wuchao Li, Xinchun Luo, Xingmei Wang, Zexuan Cheng, Zixing Zhang, Bin Zhang, Boxuan Wang, Chaoyi Ma, Chengru Song, Chenhui Wang, Di Wang, Dongxue Meng, Fan Yang, Fangyu Zhang, Feng Jiang, Fuxing Zhang, Gang Wang, Guowang Zhang, Han Li, Hengrui Hu, Hezheng Lin, Hongtao Cheng, Hongyang Cao, Huanjie Wang, Jiaming Huang, Jiapeng Chen, Jiaqiang Liu, Jinghui Jia, Kun Gai, Lantao Hu, Liang Zeng, Liao Yu, Qiang Wang, Qidong Zhou, Shengzhe Wang, Shihui He, Shuang Yang, Shujie Yang, Sui Huang, Tao Wu, Tiantian He, Tingting Gao, Wei Yuan, Xiao Liang, Xiaoxiao Xu, Xugang Liu, Yan Wang, Yi Wang, Yiwu Liu, Yue Song, Yufei Zhang, Yunfan Wu, Yunfeng Zhao, and Zhanyu Liu. 2025. OneRec Technical Report. *arXiv:2506.13695 [cs.IR]* <https://arxiv.org/abs/2506.13695>
- [16] 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. *Proceedings of the AAAI Conference on Artificial Intelligence* 33 (Jul 2019), 5941–5948. <https://doi.org/10.1609/aaai.v33i01.33015941>
- [17] Guorui Zhou, Xiaoqiang Zhu, Chenru Song, Ying Fan, Han Zhu, Xiao Ma, Yanghui Yan, Junqi Jin, Han Li, and Kun Gai. 2018. Deep Interest Network for Click-Through Rate Prediction. *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining* (Jul 2018). <https://doi.org/10.1145/3219819.3219823>