
A RELATED WORK

Our content recommendation task follows the paradigm of sequential recommendation (SRec) (Kang & McAuley, 2018; Wang et al., 2019). Accordingly, our work is closely aligned with the research on sequential recommendation and diffusion model (DM)-based sequential recommendation. In this section, we review the studies on these two topics in detail.

A.1 SEQUENTIAL RECOMMENDATION

Sequential recommendation (SRec) has been widely studied in RSs, owing to the natural temporal order of users’ behaviors (Wang et al., 2019; 2021). SRec can be technically divided into two categories: traditional sequential models and deep learning-based models. Traditional sequential models generally leverage sequential pattern mining (Yap et al., 2012) or Markov chain models (He & McAuley, 2016) to model the item dependencies in users’ interaction sequences. Traditional sequential methods can only capture simple interaction patterns or short-term dependencies, thereby cannot achieve satisfactory recommendation performance. To overcome these limitations, deep learning-based sequential recommendation methods are proposed to model complex and long-term dependencies in users’ behaviors. Among this category, one research line focuses on designing the effective sequence encoders and backbone networks to encode users’ interaction sequence, including GRU (Hidasi et al., 2015), CNNs (Tang & Wang, 2018), Transformer (Kang & McAuley, 2018), and Mamba (Liu et al., 2025b). Building upon these, another research line further introduces advanced models, such as Graph Neural Networks (GNNs) (Chang et al., 2021) and generative models (Deldjoo et al., 2024). Among them, generative models have recently attracted significant attention. In particular, DMs Liu et al. (2025a) and large language models (LLMs) (Sheng et al., 2025) have emerged as the two most prominent approaches. DM-based methods will be discussed in detail in Section A.2. LLM-based methods focus on leveraging the open-world knowledge encoded in LLMs to enhance sequential recommendation performance (Harte et al., 2023).

A.2 DIFFUSION MODELS FOR SEQUENTIAL RECOMMENDATION

In recent years, owing to the strong capability to model complex distributions of user behaviors and item content, diffusion models (DMs) have been widely applied in recommendation scenarios (Wei & Fang, 2025; Lin et al., 2024), including top-K recommendation (Wang et al., 2023b; Zhao et al., 2024) and multimodal recommendation (Ma et al., 2024c; Li et al., 2025a). In SRec, DM-based recommendation methods can be broadly categorized into two types: next item generation-based methods, and data augmentation-based methods. The former generally employ sequence encoders (e.g., GRU and Transformer) to encode users’ context items into condition embeddings, which then guide the generation of next items (Yang et al., 2023b; Liu et al., 2025a; Li et al., 2025b; Cai et al., 2025; Hu et al., 2024; Li et al., 2025b; Ma et al., 2024b; Wang et al., 2024b; Xie et al., 2024). For example, (Yang et al., 2023b; Liu et al., 2025a) utilize Transformer to learn condition embeddings from users’ historical interactions, which are then utilized to guide the next-item generation process. The latter category leverages DMs to generate additional interaction data in order to enrich users’ interaction sequences and alleviate sequence sparsity. For instance, (Liu et al., 2023; Ma et al., 2024a; Wu et al., 2023) propose generating pseudo interaction sequences with DMs to mitigate the sequence sparsity problem. Additionally, several methods integrate contrastive learning with diffusion models to generate augmented views, thereby enhancing the training of DM-based recommendation methods (Cui et al., 2024b;a; Qu & Nobuhara, 2025).

Although these methods have achieved remarkable success, they pose a significant risk of generating uncredible content recommendations (e.g., fake news (Wang et al., 2022; 2024a; Ma et al., 2025), misinformation (Pathak et al., 2023; Fernandez et al., 2024)), which can severely harm both user experience and societal well-being. While (Ma et al., 2025) attempts to leverage DMs to mitigate fake news, its effectiveness is limited under the challenge of scarce labeled data. This limitation motivates us to steer DMs towards credible content recommendation while simultaneously addressing the challenge of learning from only limited annotated data.

B MORE EXPERIMENTAL DETAILS

B.1 DATASETS

In our task setting, we require users’ chronological interaction sequences with content items, together with labels indicating whether each item contains uncredible content. However, only a limited number of public datasets fulfill these requirements. In this paper, we utilize three datasets: PolitiFact, GossipCop, and MHMisinfo.

The PolitiFact and GossipCop datasets are derived from the FakeNewsNet repository¹, which collects data from two well-known fact-checking websites: PolitiFact and GossipCop. These datasets provide user–news interaction sequences along with labels that indicate whether each news article is fake or true. The MHMisinfo dataset is collected from a video-based mental health misinformation dataset², containing users’ interaction sequences with videos annotated by whether the videos contain mental health misinformation. Although this dataset records user–video interactions, the original video and image contents are not provided. Therefore, we represent the items using their video descriptions instead of visual features.

Given the high sparsity of these datasets, we adopt a data augmentation strategy following common practice (Yang et al., 2023b;a). Specifically, for each user, we transform their interaction sequence into multiple sub-sequences by treating each item as the target item and the items preceding it as historical context. This transformation increases the number of user–item interaction sequences and enriches the training data. The statistics of these datasets are reported in Table 1. After augmentation, the datasets have more sequences, thereby the recommendation performances of Rec4Mit and HDInt are different from the results reported in (Wang et al., 2022) and (Wang et al., 2024a).

Table 1: The statistics of the three used datasets after preprocessing.

Datasets	PolitiFact	GossipCop	MHMisinfo
# Content items	616	9,529	3,160
# Credible content items	306	6,792	2,815
# Uncredible content items	310	2,737	345
# Training sequences	103,335	510,149	38,083
# Test sequences	21,490	68,002	8,060

B.2 BASELINE DESCRIPTIONS

In this section, we introduce the baseline methods used in our comparison.

Traditional sequential recommendation methods:

- **GRU4Rec** (Hidasi et al., 2015) utilizes the Gated Recurrent Unit (GRU) to model the temporal dependencies of items in users’ interaction sequences.
- **SASRec** (Kang & McAuley, 2018) employs the Transformer architecture to model the item dependencies in users’ interaction sequences. This is one of the most representative sequential recommendation methods.
- **Bert4Rec** (Sun et al., 2019) replaces SASRec’s unidirectional Transformer with a bidirectional Transformer architecture to model complex item dependencies. It also introduces a cloze task paradigm for sequential recommendation.
- **LRU4Rec** (Yue et al., 2024) designs linear recurrent units for sequential recommendation. It decomposes linear recurrence operations and proposes recursive parallelization, reducing model size and enabling efficient parallel training.

Contrastive learning-based sequential recommendation methods:

- **CL4SRec** (Xie et al., 2022) uses contrastive learning to address the data sparsity problem in sequential recommendation. It designs three sequence augmentation operations for contrastive

¹<https://github.com/KaiDMML/FakeNewsNet>

²<https://zenodo.org/records/13191247>

learning: item cropping, item masking, and item reordering. Transformer is used as the sequential encoder of CL4SRec.

- **ContraRec** (Wang et al., 2023a) proposes two types of contrastive perspectives to enhance the performance of contrastive learning-based sequential recommendation: context-target contrast and context-context contrast. Transformer is used as the sequential encoder of ContraRec.

Sequential recommendation methods for mitigating incredible content:

- **Rec4Mit** (Wang et al., 2022) first utilizes a disentangler to extract event- and veracity-aware information, respectively. Thereafter, the event embeddings are utilized to derive users’ genuine preferences and predict the next items users may be interested in.
- **HDInt** (Wang et al., 2024a). Similar to Rec4Mit, HDInt is also dedicated to mitigating fake news in recommender systems. HDInt also considers the political bias. We omit this part, since it requires additional data and the political bias is not considered in our task.
- **PRISM** (Ma et al., 2025) proposes a protection-enhanced news recommendation method based on interest-aware sequential modeling. It utilizes DMs’ controllable ability to learn user interest and mitigate fake news. However, it assumes all the labels of fake news are fully available, which does not hold in the real world. It is also a DM-based sequential recommendation method.

DM-based recommendation methods:

- **DreamRec** (Yang et al., 2023b) assumes that each user has an “oracle” item in mind and selects items that match his ideal item. It uses a Transformer to learn users’ preferences, which then serve as the condition for generating the oracle item for each user.
- **DiffuRec** (Li et al., 2023) employs a diffusion model to represent item embeddings in a distribution space and then feeds the embeddings into an approximator to generate target item representations. It argues that the standard objective function of DMs is unsuitable for recommendation tasks and uses cross-entropy loss to optimize model parameters.
- **PreferDiff** (Liu et al., 2025a) proposes a surrogate optimization objective which extend BPR recommendation loss (Rendle et al., 2009) to variational format. Meanwhile, this surrogate optimization objective can also be extended to multiple negative items.

B.3 EVALUATION METRICS

HR@K and NDCG@K are two commonly used metrics to evaluate the recommendation accuracy, thereby we do not make further introduction for them. Credible Rate (CR@K) is a metric to measure the credibility of a recommendation model. Specifically, it calculates the average rate of the credible content items in the recommendation lists:

$$CR@K = \frac{1}{|\mathcal{S}_{test}|} \sum_{s \in \mathcal{S}_{test}} \frac{K - |\mathcal{R}_s \cap \mathcal{I}_{unc}^{Ground-truth}|}{K}, \quad (1)$$

where $\mathcal{S}_{test}$ is the test set of sequences. $\mathcal{R}_s$ is the recommendation list for sequence s . $\mathcal{I}_{unc}^{Ground-truth}$ denotes the ground-truth set of incredible items. $|\mathcal{R}_s \cap \mathcal{I}_{unc}^{Ground-truth}|$ calculates the number of incredible items in the recommendation list. The higher value of CR@K means the better performance in delivering credible recommendations.

In addition, we test how our methods perform in terms of both accurate and credible recommendations, we design a combined metric HC@K (i.e., combining HR@K and CR@K). Formally, HC@K is calculated as follows:

$$HC@K = \frac{2 \times HR@K \times (CR@K/2)}{HR@K + (CR@K/2)}. \quad (2)$$

This combined metric is inspired by the F1-score, which combines precision rate and recall rate. To note that, since the values of HR@K and CR@K are not on the same scale, we divide CR@K with a factor of 2 to rescale it into a similar value level with HR@K. This adjustment ensures a fair combination; otherwise, the metric with a much smaller magnitude would disproportionately dominate the combined score.

B.4 IMPLEMENTATION DETAILS

In this paper, we consider a more challenging and realistic scenario in which only a small proportion of uncredible items are verified. To simulate this setting, we randomly select 20% of the uncredible items with available labels during the training process, while the labels of the remaining items are treated as unknown. It is similar to the semi-supervised setting. In contrast, during the testing stage, all content labels are provided to enable an accurate evaluation.

The items in PolitiFact and GossipCop are news articles, and we use their textual descriptions as item content. In MHMisinfo, although the items are videos, only textual descriptions are available; thus, we can only rely on the textual descriptions for content representation. We encode these textual descriptions into language embeddings using LLaMA2-7B (Touvron et al., 2023), and further project them into a lower dimension through an MLP. Following (Liu et al., 2025a), we fix the transformed embedding dimension at 3072 for all DM-based methods, as they exhibit strong performance only with higher embedding sizes. For other methods, the embedding size is set to 64. We also experimented with larger embedding sizes for these methods, but observed little or no performance gain, and even performance drops for some methods, consistent with the findings in (Liu et al., 2025a).

In our implementation, we select Transformer as our sequence encoder. Following the standard configuration (Vaswani et al., 2017), the Transformer architecture in our implementation includes multi-head attention, position-wise feed-forward network, layer normalization, and dropout.

For our method `Disco`, the hyperparameter w is tuned within $\{0.5, 1, 1.5, 2, 5\}$. We fix m at 10,000 and tune γ within $\{0.1, 0.2, 0.3, 0.4, 0.5\}$ to control the maximum selection ratio as well as the growth rate of the current selection ratio. The maximum number of diffusion steps is fixed at 2,000 and the DDIM step is set to 100, following the settings of (Liu et al., 2025a). For all DM-based methods, we utilize a linear schedule for β_t in range $[0.0001, 0.02]$. In our implementation, we do not use a classifier-free guidance (Ho & Salimans), since we found it does not influence much to the performance of `Disco`. In our implementation, we found that the singular values in Λ are relatively large; therefore, the threshold for constructing the null space of uncredible features is fixed at 3 for all datasets in our experiments. We search learning rate in range $\{1e-5, 5e-5, 1e-4, 5e-4, 1e-3\}$. The batch size is searched in $\{2048 \times 2^i\}_{i=0,1,2,3}$. The model parameters are initialized using normal initialization and optimized by AdamW (Loshchilov & Hutter, 2017). The hyperparameter settings of baseline methods are reported in Table 2. All experiments are conducted on an NVIDIA A40 GPU with 48 GB of memory. Each method is run five times, and we report the average performance along with the standard deviation.

Table 2: The hyperparameter settings of baseline methods.

Methods	Hyperparameter searching space
GRU4Rec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0
SASRec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0
Bert4Rec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, mask probability $\sim \{0.2, 0.4, 0.6, 0.8\}$
LRURec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, dropout rate $\sim \{0.2, 0.4, 0.6, 0.8\}$
CL4SRec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, mask/reorder/crop proportion $\sim \{0.2, 0.4, 0.6, 0.8\}$, $\lambda \sim \{0.1, 0.3, \dots, 0.9\}$
ContraRec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, mask/reorder/crop proportion $\sim \{0.2, 0.4, 0.6, 0.8\}$, $\tau_1, \tau_2 \sim \{0.1, 0.2, \dots, 1\}$, $\gamma \sim \{0, 0.01, 0.1, 1, 5, 10\}$
Rec4Mit	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, $k \sim \{2, 4, \dots, 20\}$
HDInt	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, $\lambda \sim \{1, 2, \dots, 10\}$, $\gamma \sim \{2, 4, 6, 8, 10\}$
PRISM	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, $T \sim \{500, 1000, 1500, 2000\}$, $w \sim \{0.2, 0.4, 0.6, 0.8\}$, $\lambda_{OT}, \lambda_c, \lambda_r, \lambda_{rec} \sim \{0.2, 0.4, 0.6, 0.8, 1\}$, embedding size $\sim \{64, 128, 256, 512, 1024, 2048, 3072\}$
DreamRec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, $T \sim \{500, 1000, 1500, 2000\}$, $w \sim \{0.2, 0.4, 0.6, 0.8\}$, embedding size $\sim \{64, 128, 256, 512, 1024, 2048, 3072\}$
DiffuRec	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, $T \sim \{16, 32, 64, 128\}$, $\delta=0.001$, embedding size $\sim \{64, 128, 256, 512, 1024, 2048, 3072\}$
PreferDiff	$\text{lr} \sim \{1e-2, 5e-2, 1e-3, 5e-3, 1e-4\}$, weight decay=0, $T \sim \{500, 1000, 1500, 2000\}$, $w \sim \{0.2, 0.4, 0.6, 0.8\}$, $\lambda \sim \{0.2, 0.4, 0.6, 0.8\}$, embedding size $\sim \{64, 128, 256, 512, 1024, 2048, 3072\}$

B.5 MORE HYPERPARAMETER EXPERIMENTS

The hyperparameter γ controls the selection ratio of potential uncredible items. We evaluate the performance of `Disco` (using combined metric HC@5) under different values of γ in range $\{0.1, 0.2, 0.3, 0.4, 0.5\}$. As shown in Figure 1, `Disco` achieves the best performance when fixing $w = 0.1$ on PolitiFact and GossipCop, and $w = 0.4$ on MHMisinfo. Lower values prevent the model

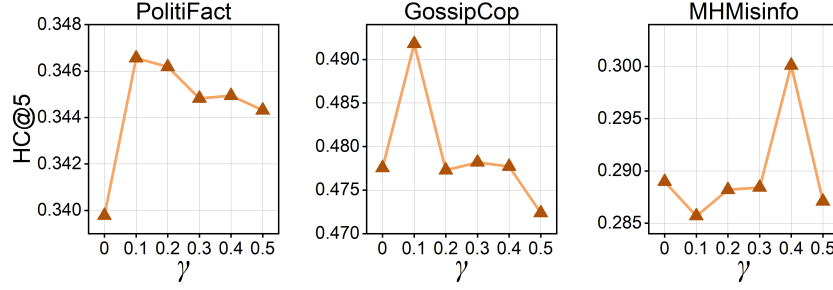


Figure 1: Effect of γ on Disco .

from effectively capturing potentially uncredible items, while higher values may introduce excessive noise, both of which degrade model performance.

B.6 TIME EFFICIENCY ANALYSIS

We conduct experiments to evaluate the training and inference cost (in seconds) of our model Disco and four DM-based methods under the same batch size. As shown in Table 3, the training cost of Disco is relatively higher than DreamRec and PreferDiff, mainly due to our additional designs for credible content recommendation, including content disentanglement and credible subspace projection. This is acceptable due to the higher recommendation accuracy and credibility of our proposed method. Our training cost is much lower than that of DiffuRec and PRISM. As for inference cost, our proposed method Disco demonstrates the highest efficiency. This is because we adopt DDIM (Song et al., 2021) as our generation strategy, which is more efficient than the DDPM (Ho et al., 2020) paradigm employed by DreamRec, DiffuRec and PRISM. Even compared with PreferDiff, which also adopts DDIM, Disco also exhibits higher efficiency. It is because we do not employ classifier-free guidance in our implementation, since it has limited influence on our model while incurring additional time consumption. Although Disco requires disentangling item embeddings first in the inference stage, its inference cost remains comparable to PreferDiff on GossipCop dataset, which contain large number of items.

Table 3: Time cost (s) of different models on PolitiFact, GossipCop, and MHMisinfo.

Datasets	Time cost (s)	DreamRec	DiffuRec	PRISM	PreferDiff	Disco
PolitiFact	Training/epoch	9.4	74.8	25.3	11.3	12.6
	Inference	278.2	224.3	994.3	15.2	10.5
GossipCop	Training/epoch	43.3	376.9	137.6	58.8	73.1
	Inference	956.8	783.1	3223.6	127.9	133.5
MHMisinfo	Training/epoch	3.3	26.9	8.9	3.9	4.2
	Inference	107.6	87.1	372.9	8.7	7.7

B.7 DISCUSSION ON EMBEDDING DIMENSION

As pointed out in our paper, DM-based methods can only achieve strong performance when the embedding dimension is high. To demonstrate the necessity of using high embedding dimension (i.e., 3072) for all DM-based methods, we conducted experiments on DM-based methods when fixing embedding dimension to 64. The results are shown in Table 4:

As shown in Table 4, not only Disco but also other DM-based recommendation methods (PRISM, DreamRec, and PreferDiff) experience a substantial performance drop when the embedding size is reduced to 64. This observation is consistent with the findings reported in [1] and further validates the necessity of using high-dimensional embeddings in DM-based recommender systems.

Although a higher embedding dimension increases the per-epoch training time, it also provides the benefit of significantly faster convergence. In our revised manuscript, we include a figure illustrating the convergence curves of Disco and SASRec. As presented in Figure 2, Disco converges much

Table 4: Performance comparison under embedding dimension 64 for DM-based methods.

Methods	HR@5	HR@10	NDCG@5	NDCG@10	CR@5	CR@10	HC@5	HC@10
PolitiFact								
GRU4Rec	0.2142	0.3390	0.1463	0.1863	0.9266	0.9122	0.2929	0.3889
SASRec	0.2158	0.3519	0.1386	0.1823	0.9059	0.9028	0.2923	0.3955
Bert4Rec	0.2191	0.3473	0.1472	0.1883	0.9127	0.9045	0.2960	0.3929
CL4SRec	0.2247	0.3527	0.1508	0.1919	0.9132	0.9027	0.3012	0.3960
PRISM (emb_size=3072)	0.1927	0.2758	0.1348	0.1615	0.9325	0.9178	0.2727	0.3446
PRISM (emb_size=64)	0.0806	0.1261	0.0569	0.0715	0.7807	0.7222	0.1336	0.1869
DreamRec (emb_size=3072)	0.2416	0.3287	0.1767	0.2047	0.8620	0.8437	0.3054	0.3661
DreamRec (emb_size=64)	0.0814	0.1074	0.0651	0.0734	0.5744	0.5605	0.1268	0.1553
PreferDiff (emb_size=3072)	0.2531	0.3554	0.1818	0.2147	0.8925	0.8981	0.3228	0.3968
PreferDiff (emb_size=64)	0.1227	0.1841	0.0882	0.1078	0.8934	0.8788	0.1925	0.2595
Disco (emb_size=3072)	0.2678	0.3775	0.1983	0.2336	0.9823	0.9425	0.3466	0.4192
Disco (emb_size=64)	0.1171	0.1962	0.1107	0.1422	0.9916	0.9665	0.1895	0.2791
GossipCop								
GRU4Rec	0.2226	0.3194	0.1466	0.1778	0.8864	0.8706	0.2957	0.3678
SASRec	0.3078	0.4706	0.1607	0.2135	0.8743	0.8526	0.3612	0.4473
Bert4Rec	0.2372	0.3711	0.1338	0.1770	0.8764	0.8587	0.3078	0.3981
CL4SRec	0.2898	0.4100	0.1784	0.2174	0.8938	0.8932	0.3516	0.4275
PRISM (emb_size=3072)	0.2948	0.3447	0.2301	0.2463	0.8806	0.8733	0.3531	0.3852
PRISM (emb_size=64)	0.0023	0.0034	0.0015	0.0018	0.6570	0.6940	0.0046	0.0067
DreamRec (emb_size=3072)	0.4619	0.5501	0.3415	0.3704	0.8464	0.8336	0.4415	0.4742
DreamRec (emb_size=64)	0.0036	0.0049	0.0027	0.0031	0.5791	0.5903	0.0071	0.0096
PreferDiff (emb_size=3072)	0.4969	0.6022	0.3655	0.3999	0.8307	0.8228	0.4523	0.4887
PreferDiff (emb_size=64)	0.0084	0.0139	0.0053	0.0070	0.6542	0.7400	0.0164	0.0268
Disco (emb_size=3072)	0.5236	0.6143	0.3996	0.4292	0.9272	0.9039	0.4918	0.5207
Disco (emb_size=64)	0.0087	0.0162	0.0053	0.0077	0.7993	0.7993	0.0170	0.0311
MHMisinfo								
GRU4Rec	0.1151	0.1894	0.0760	0.0998	0.8380	0.8608	0.1803	0.2624
SASRec	0.1485	0.2592	0.0826	0.1179	0.8339	0.8915	0.2190	0.3276
Bert4Rec	0.1391	0.2299	0.0847	0.1138	0.8162	0.8786	0.2074	0.3017
CL4SRec	0.1734	0.2621	0.1101	0.1387	0.8577	0.9081	0.2469	0.3323
PRISM (emb_size=3072)	0.1700	0.2339	0.1181	0.1388	0.8398	0.8919	0.2418	0.3065
PRISM (emb_size=64)	0.0239	0.0295	0.0190	0.0208	0.7095	0.7317	0.0448	0.0546
DreamRec (emb_size=3072)	0.1819	0.2426	0.1313	0.1509	0.9002	0.8952	0.2633	0.3176
DreamRec (emb_size=64)	0.0282	0.0347	0.0233	0.0254	0.8281	0.8901	0.0528	0.0644
PreferDiff (emb_size=3072)	0.1974	0.2620	0.1389	0.1598	0.8693	0.8874	0.2713	0.3294
PreferDiff (emb_size=64)	0.0325	0.0380	0.0290	0.0308	0.8290	0.8989	0.0603	0.0701
Disco (emb_size=3072)	0.2215	0.2822	0.1580	0.1778	0.9305	0.9264	0.3000	0.3507
Disco (emb_size=64)	0.0123	0.0572	0.0074	0.0213	0.9968	0.9873	0.0240	0.1025

more rapidly than the non-DM-based method SASRec (embedding size = 64). Specifically, Disco reaches its best performance at approximately the 40-th epoch, whereas SASRec requires around 400 epochs. This fast convergence rate of Disco partially offsets the additional computational cost introduced by high-dimensional embeddings.

In addition, we conducted further experiments to demonstrate that Disco can still achieve superior performance compared with non-DM-based methods when the embedding dimension is restricted to 64, as long as a minor modification is applied to the overall optimization objective. In particular, we augment Disco’s loss function with an additional Cross-Entropy term:

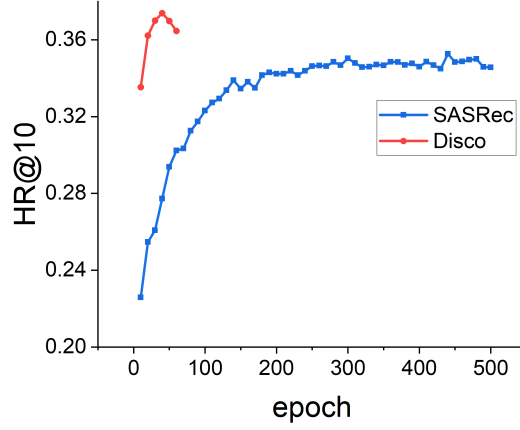


Figure 2: The performance convergence curve of SASRec and Disco on PolitiFact dataset.

Table 5: Performance comparison of Disco (optimized by $\mathcal{L}_{\text{Disco}^*}$) and non-DM recommendation methods on PolitiFact, GossipCop, and MHMisinfo when embedding dimension is set to 64.

Methods	HR@5	HR@10	NDCG@5	NDCG@10	CR@5	CR@10	HC@5	HC@10
PolitiFact								
GRU4Rec	0.2142	0.3390	0.1463	0.1863	0.9266	0.9122	0.2929	0.3889
SASRec	0.2158	0.3519	0.1386	0.1823	0.9059	0.9028	0.2923	0.3955
Bert4Rec	0.2191	0.3473	0.1472	0.1883	0.9127	0.9045	0.2960	0.3929
CL4SRec	0.2247	0.3527	0.1508	0.1919	0.9132	0.9027	0.3012	0.3960
Disco (emb_size=64, $\mathcal{L}_{\text{Disco}^*}$)	0.2335	0.3555	0.1642	0.2034	0.9316	0.9213	0.3111	0.4013
GossipCop								
GRU4Rec	0.2226	0.3194	0.1466	0.1778	0.8864	0.8706	0.2957	0.3678
SASRec	0.3078	0.4706	0.1607	0.2135	0.8743	0.8526	0.3612	0.4473
Bert4Rec	0.2372	0.3711	0.1338	0.1770	0.8764	0.8587	0.3078	0.3981
CL4SRec	0.2898	0.4100	0.1784	0.2174	0.8938	0.8932	0.3516	0.4275
Disco (emb_size=64, $\mathcal{L}_{\text{Disco}^*}$)	0.4250	0.5060	0.3320	0.3583	0.9151	0.9080	0.4407	0.4786
MHMisinfo								
GRU4Rec	0.1151	0.1894	0.0760	0.0998	0.8380	0.8608	0.1803	0.2624
SASRec	0.1485	0.2592	0.0826	0.1179	0.8339	0.8915	0.2190	0.3276
Bert4Rec	0.1391	0.2299	0.0847	0.1138	0.8162	0.8786	0.2074	0.3017
CL4SRec	0.1734	0.2621	0.1101	0.1387	0.8577	0.9081	0.2469	0.3323
Disco (emb_size=64, $\mathcal{L}_{\text{Disco}^*}$)	0.1856	0.2660	0.1214	0.1475	0.8669	0.9129	0.2599	0.3361

$$\mathcal{L}_{\text{Disco}^*} = \mathcal{L}_{\text{Disco}} - \log \left(\frac{\exp(f_{\theta}(\tilde{\mathbf{e}}_n^t, \mathbf{c}^{\text{pre}}, t) \cdot \mathbf{e}_n^{\top})}{\sum_{i \in \mathcal{I}} \exp(f_{\theta}(\tilde{\mathbf{e}}_n^t, \mathbf{c}^{\text{pre}}, t) \cdot \mathbf{e}_i^{\top})} \right). \quad (3)$$

The added term encourages the diffusion network f_{θ} to align its predictions more closely to the target items than with other items. Using this enhanced loss function $\mathcal{L}_{\text{Disco}^*}$, Disco can achieve better performance than non-DM based methods. The comparison results are reported in Table 5.

As shown in the Table 5, Disco can achieve better performance than non-DM recommendation methods with only a minor adjustment to its overall loss. This improvement arises because adding a discriminative loss (i.e., Cross Entropy loss) to the generative loss ($\mathcal{L}_{\text{Disco}}$) can partially mitigate the dimensionality limitations inherent to diffusion models.

However, this practice will transform a purely generative model into a discriminative one. Our work does not make such a compromise. Nevertheless, the results clearly suggest that our model retains strong potential to surpass non-DM-based recommenders even when operating with substantially reduced embedding dimensions. This highlights the potential of `Disco` to achieve an effective trade-off between efficiency and performance.

B.8 DISCUSSION ON DIFFUSION STEP T

In this section, we conducted experiments to investigate whether `Disco` still strong performance when the diffusion step T is much smaller. The experimental results are reported in Table 6.

From Table 6, we can have the following observations:

- As the diffusion step T increases, the performance of our proposed model, `Disco`, improves.
- `Disco` maintains superior performance compared with non-DM methods even at much smaller diffusion step ($T=100$ for `PolitiFact` and `GossipCop`, and $T=500$ for `MHMisinfo`).

Table 6: Performance comparison of `Disco` and non-DM recommendation methods under different diffusion steps T .

Methods	HR@5	HR@10	NDCG@5	NDCG@10	CR@5	CR@10	HC@5	HC@10
PolitiFact								
GRU4Rec	0.2142	0.3390	0.1463	0.1863	0.9266	0.9122	0.2929	0.3889
SASRec	0.2158	0.3519	0.1386	0.1823	0.9059	0.9028	0.2923	0.3955
Bert4Rec	0.2191	0.3473	0.1472	0.1883	0.9127	0.9045	0.2960	0.3929
CL4SRec	0.2247	0.3527	0.1508	0.1919	0.9132	0.9027	0.3012	0.3960
<code>Disco</code> (Diffusion step $T=100$)	0.2494	0.3724	0.1751	0.2146	0.9488	0.9352	0.3269	0.4146
<code>Disco</code> (Diffusion step $T=200$)	0.2602	0.3752	0.1842	0.2211	0.9427	0.9340	0.3353	0.4161
<code>Disco</code> (Diffusion step $T=500$)	0.2591	0.3828	0.1811	0.2208	0.9434	0.9272	0.3345	0.4193
<code>Disco</code> (Diffusion step $T=1000$)	0.2525	0.3784	0.1752	0.2156	0.9488	0.9369	0.3296	0.4186
<code>Disco</code> (Diffusion step $T=2000$)	0.2678	0.3775	0.1983	0.2336	0.9823	0.9425	0.3466	0.4192
GossipCop								
GRU4Rec	0.2226	0.3194	0.1466	0.1778	0.8864	0.8706	0.2957	0.3678
SASRec	0.3078	0.4706	0.1607	0.2135	0.8743	0.8526	0.3612	0.4473
Bert4Rec	0.2372	0.3711	0.1338	0.1770	0.8764	0.8587	0.3078	0.3981
CL4SRec	0.2898	0.4100	0.1784	0.2174	0.8938	0.8932	0.3516	0.4275
<code>Disco</code> (Diffusion step $T=100$)	0.4603	0.5479	0.3419	0.3705	0.9304	0.9252	0.4627	0.5016
<code>Disco</code> (Diffusion step $T=200$)	0.4759	0.5659	0.3537	0.3831	0.9242	0.9199	0.4689	0.5075
<code>Disco</code> (Diffusion step $T=500$)	0.4867	0.5798	0.3621	0.3925	0.9258	0.9169	0.4745	0.5120
<code>Disco</code> (Diffusion step $T=1000$)	0.4936	0.5956	0.3651	0.3984	0.9202	0.9039	0.4763	0.5139
<code>Disco</code> (Diffusion step $T=2000$)	0.5236	0.6143	0.3996	0.4292	0.9272	0.9039	0.4918	0.5207
MHMisinfo								
GRU4Rec	0.1151	0.1894	0.0760	0.0998	0.8380	0.8608	0.1803	0.2624
SASRec	0.1485	0.2592	0.0826	0.1179	0.8339	0.8915	0.2190	0.3276
Bert4Rec	0.1391	0.2299	0.0847	0.1138	0.8162	0.8786	0.2074	0.3017
CL4SRec	0.1734	0.2621	0.1101	0.1387	0.8577	0.9081	0.2469	0.3323
<code>Disco</code> (Diffusion step $T=100$)	0.1378	0.1998	0.0914	0.1111	0.8526	0.8783	0.2083	0.2746
<code>Disco</code> (Diffusion step $T=200$)	0.1393	0.2191	0.0921	0.1178	0.9161	0.9209	0.2136	0.2969
<code>Disco</code> (Diffusion step $T=500$)	0.1819	0.2610	0.1299	0.1553	0.9076	0.9144	0.2597	0.3323
<code>Disco</code> (Diffusion step $T=1000$)	0.2144	0.2638	0.1547	0.1708	0.9379	0.9311	0.2943	0.3358
<code>Disco</code> (Diffusion step $T=2000$)	0.2215	0.2822	0.1580	0.1778	0.9305	0.9264	0.3000	0.3507

B.9 DISCUSSION ON VARIOUS RATIOS OF AVAILABLE CREDIBILITY LABELS

we conduct additional experiments under different credibility label ratio of uncredible content (i.e., 5%, 10%, 20%, 30%, 50%). The experimental results are reported in Table 7. From the results reported in the Table, we can have the following findings:

- **Finding 1:** As the credibility label ratio increases, the recommendation credibility (CR) improves steadily. This is because a larger number of credibility labels enables the construction of a more comprehensive and accurate credible subspace, allowing uncredible content to be mitigated more effectively.
- **Finding 2:** The recommendation accuracy fluctuates only slightly within a narrow range across different label ratios. This stability is attributed to our proposed disentangled diffusion model, which effectively mitigates uncredible content while preserving users' preference-related information, thereby maintaining high recommendation accuracy.

Table 7: Performance comparison under different credibility label ratios.

Label ratios	HR@5	HR@10	NDCG@5	NDCG@10	CR@5	CR@10	HC@5	HC@10
PolitiFact								
5%	0.2617	0.3869	0.1819	0.2222	0.9422	0.9279	0.3365	0.4219
10%	0.2541	0.3836	0.1768	0.2184	0.9476	0.9357	0.3308	0.4216
20%	0.2678	0.3775	0.1983	0.2336	0.9823	0.9425	0.3466	0.4192
30%	0.2704	0.3838	0.1980	0.2345	0.9829	0.9518	0.3489	0.4249
50%	0.2678	0.3832	0.1923	0.2294	0.9842	0.9486	0.3468	0.4239
GossipCop								
5%	0.5179	0.6021	0.4003	0.4279	0.9261	0.8726	0.4889	0.5060
10%	0.5290	0.6115	0.4089	0.4359	0.9266	0.8764	0.4940	0.5105
20%	0.5236	0.6143	0.3996	0.4292	0.9272	0.9039	0.4918	0.5207
30%	0.5196	0.6141	0.3947	0.4255	0.9278	0.9101	0.4902	0.5227
50%	0.5151	0.6069	0.3953	0.4253	0.9284	0.9176	0.4883	0.5226
MHMisinfo								
5%	0.2127	0.2762	0.1500	0.1705	0.9149	0.9114	0.2904	0.3439
10%	0.2112	0.2798	0.1506	0.1728	0.9217	0.9152	0.2897	0.3473
20%	0.2215	0.2822	0.1580	0.1778	0.9305	0.9264	0.3001	0.3507
30%	0.2207	0.2836	0.1551	0.1755	0.9303	0.9244	0.2994	0.3515
50%	0.2134	0.2715	0.1533	0.1721	0.9331	0.9279	0.2929	0.3425

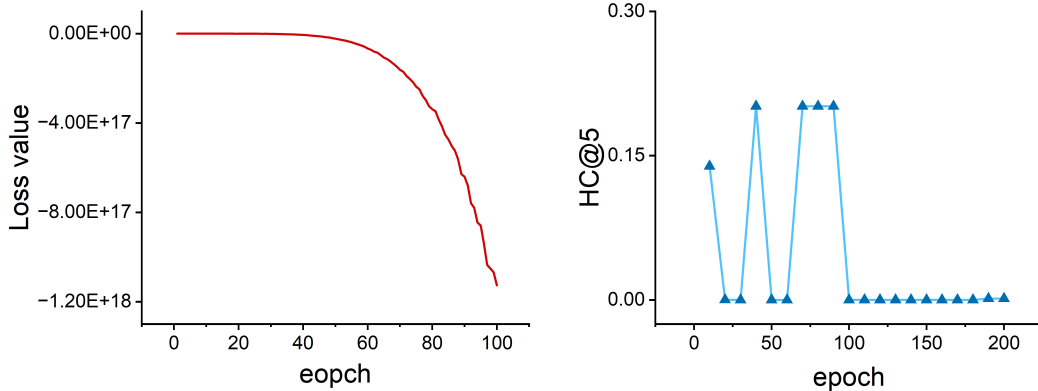


Figure 3: The loss and performance (HC@5) curves of variant w/o CE on PolitiFact dataset.

B.10 DISCUSSION ON THE INSTABILITY OF ABLATION VARIANT W/O CE

In this section, we conduct an empirical study on the convergence of the ablation variant "w/o CE". As show in Figure 3, we observe that the variant w/o CE (i.e., not replacing the MSE error with cosine error) leads to extremely unstable training and performance. Specifically, the loss rapidly collapses to an extremely small value (around -1.2×10^{18}), and the performance (HR@5) exhibits severe fluctuations. These results verify the necessity of replacing the MSE error with the cosine error to ensure stable optimization.

C WHY DMS POSE A DANGER OF GENERATING UNCREDIBLE CONTENT RECOMMENDATION?

In this section, we empirically and theoretically analyze why existing DM-based recommendation methods risk generating incredible recommendations.

C.1 EMPIRICAL FINDINGS

In DM-based recommendation methods, the condition and diffusion target are two critical factors. In this section, we conduct experiments to examine how they influence the recommendation credibility of DM-based methods. Specifically, we divide the training dataset into four subsets based on whether the context items or the diffusion target (i.e., target items) contain incredible content. We use $\checkmark$ to denote that context items or target items contain incredible content, and $\times$ to denote the opposite. After this dataset division, we train two representative DM-based recommendation methods (DreamRec (Yang et al., 2023b) and PreferDiff (Liu et al., 2025a)) on each subset. From the results reported in Table 8, we can find that these two factors indeed affect the recommendation credibility of DM-based methods. We refer to these two factors as incredible condition and incredible diffusion target.

- **Uncredible condition.** When controlling the diffusion target, if the context items contain uncredible content that leads to an uncredible condition, the credibility metric CR@10 (i.e., credible Rate) decreases to some extent for both DreamRec and PreferDiff across the PolitiFact and GossipCop datasets. This finding indicates that the uncredible condition is a factor contributing to the risk of DMs generating uncredible recommendation results.
- **Uncredible diffusion target.** When controlling the condition, if the diffusion target is an uncredible item (i.e., an uncredible diffusion target), CR@10 drops to an extremely low level. This further emphasizes that the uncredible diffusion target is another key contributing factor.

Apart from these two findings, we also observe that training with the complete datasets yields worse recommendation credibility compared to the subset where neither the condition nor the diffusion target contains uncredible items. This further validates that the uncredible condition and the uncredible diffusion target are indeed the key contributing factors that place DM-based recommendation methods at risk of generating uncredible recommendation results.

Moreover, **although simply removing uncredible items from the datasets can improve recommendation credibility, it significantly deteriorates recommendation accuracy.** This is because uncredible items may also reflect users' genuine preferences, thereby discarding them restricts the model's ability to accurately learn users' true interests. Therefore, it is crucial to design advanced models that can mitigate the recommendation of uncredible content while simultaneously preserving high recommendation accuracy. This is the motivation and research significance of our proposed model, Disco.

C.2 THEORETICAL ANALYSIS

Proof: Uncredible condition can enhance DM's generation of uncredible results

The training of a conditional DM is to maximize $\mathbb{E}_{p_{data}(\mathbf{e}_n, \mathbf{c})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c})]$, where $\mathbf{e}_n$ is the diffusion target (i.e., the last item in a user's interaction sequence) and $\mathbf{c}$ is the condition. This training objective pushes the generation toward the real data distribution.

Table 8: Performance comparison of DreamRec and PreferDiff under different settings of unbelievable content items in condition and diffusion target on PolitiFact and GossipCop datasets. Best results are highlighted in bold.

Methods	Whether contain unbelievable content items?		Politi				Gossip			
	condition	diffusion target	HR@10	NDCG@10	CR@10	HC@10	HR@10	NDCG@10	CR@10	HC@10
DreamRec	Training with complete dataset		0.3287	0.2047	0.8437	0.3661	0.5501	0.3704	0.8336	<u>0.4742</u>
	$\times$	$\times$	0.2674	0.1571	0.9935	0.3477	0.4658	0.3160	0.9771	0.4769
	$\times$	$\checkmark$	0.0577	0.0409	0.1888	0.0716	0.0372	0.0284	0.0522	0.0307
	$\checkmark$	$\times$	0.2671	0.1541	<u>0.9875</u>	0.3467	0.1927	0.1368	<u>0.9340</u>	0.2728
	$\checkmark$	$\checkmark$	0.0684	0.0413	0.0806	0.0507	0.0539	0.0404	0.0450	0.0317
PreferDiff	Training with complete dataset		0.3554	0.2147	<u>0.8981</u>	0.3968	0.6022	0.3999	0.8228	0.4887
	$\times$	$\times$	0.3035	0.1915	0.9591	0.3717	0.5036	0.3657	0.9315	<u>0.4839</u>
	$\times$	$\checkmark$	0.0557	0.0410	0.1073	0.0547	0.0407	0.0304	0.0833	0.0412
	$\checkmark$	$\times$	0.2625	0.1553	0.8561	0.3254	0.2074	0.1454	<u>0.9076</u>	0.2847
	$\checkmark$	$\checkmark$	0.0568	0.0385	0.0837	0.0482	0.0573	0.0421	0.0254	0.0208

When an unbelievable content-related condition $\mathbf{c}^{unc}$ is utilized to guide the generation process, the model aims to maximize $\mathbb{E}_{p_{data}(\mathbf{e}_n, \mathbf{c}^{unc})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc})]$. Then, we have:

$$\begin{aligned}
\mathbb{E}_{p_{data}(\mathbf{e}_n, \mathbf{c}^{unc})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc})] &= \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} \mathbb{E}_{p_{data}(\mathbf{e}_n | \mathbf{c}^{unc})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc})] \\
&= \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} \left[\int_{\mathbf{e}_n} p_{data}(\mathbf{e}_n | \mathbf{c}^{unc}) \log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc}) d\mathbf{e}_n \right] \\
&= \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} [-H(p_{data}(\mathbf{e}_n^* | \mathbf{c}^{unc}), p_{\theta}(\mathbf{e}_n^* | \mathbf{c}^{unc}))] \\
&= \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} [-H(p_{data}(\mathbf{e}_n^* | \mathbf{c}^{unc}))] \\
&\quad - \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} [D_{KL}(p_{data}(\mathbf{e}_n^* | \mathbf{c}^{unc}) || p_{\theta}(\mathbf{e}_n^* | \mathbf{c}^{unc}))] \\
&= -H_{p_{data}}(\mathcal{E} | \mathcal{C}^{unc}) \\
&\quad - \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} [D_{KL}(p_{data}(\mathbf{e}_n^* | \mathbf{c}^{unc}) || p_{\theta}(\mathbf{e}_n^* | \mathbf{c}^{unc}))],
\end{aligned} \tag{4}$$

where $\mathcal{E}$ represents the whole generation space and $\mathbf{e}_n^* \in \mathcal{E}$. $\mathcal{C}^{unc}$ is the whole space of unbelievable condition $\mathbf{c}^{unc}$. $H(\cdot, \cdot)$ is the entropy between two variables or distributions. According to the above derivation, we have:

$$\begin{aligned}
H_{p_{data}}(\mathcal{E} | \mathcal{C}^{unc}) &= -\mathbb{E}_{p_{data}(\mathbf{e}_n, \mathbf{c}^{unc})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc})] \\
&\quad - \mathbb{E}_{p_{data}(\mathbf{c}^{unc})} [D_{KL}(p_{data}(\mathbf{e}_n^* | \mathbf{c}^{unc}) || p_{\theta}(\mathbf{e}_n^* | \mathbf{c}^{unc}))] \\
&\leq -\mathbb{E}_{p_{data}(\mathbf{e}_n, \mathbf{c}^{unc})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc})].
\end{aligned} \tag{5}$$

Ideally, when the model is optimally trained, the D_{KL} term will approach zero, indicating that the conditional generation distribution approaches the real data distribution. Therefore, the mutual information between the whole conditional generation space $\mathcal{E}$ and the whole unbelievable condition space $\mathcal{C}^{unc}$ can be calculated as:

$$\begin{aligned}
I_{p_{\theta}}(\mathcal{E}, \mathcal{C}^{unc}) &= I_{p_{data}}(\mathcal{E}, \mathcal{C}^{unc}) \\
&= H_{p_{data}}(\mathcal{E}) - H_{p_{data}}(\mathcal{E} | \mathcal{C}^{unc}) \\
&\geq H_{p_{data}}(\mathcal{E}) + \mathbb{E}_{p_{data}(\mathbf{e}_n, \mathbf{c}^{unc})} [\log p_{\theta}(\mathbf{e}_n | \mathbf{c}^{unc})].
\end{aligned} \tag{6}$$

The second equation is derived according to the property of mutual information: $I(X, Y) = H(X) - H(X|Y)$. As the training goes on, the second term becomes larger. At the same time, $H_{p_{data}}(\mathcal{E})$ is a constant based on the real data distribution p_{data} . Hence, the lower bound of $I_{p_{\theta}}(\mathcal{E}, \mathcal{C}^{unc})$ also becomes larger. Based on this, we can conclude that training the diffusion model with unbelievable conditions increases the mutual information between the generation space and the unbelievable condition space. This indicates that the generation space increasingly contains unbelievable features reflected in the unbelievable conditions.

Proof: Unbelievable diffusion target can enhance DM's generation of unbelievable results

The optimization loss of existing DM-based recommendation methods can be formulated as:

$$\mathcal{L} = \mathbb{E}_{t \sim U(0, T)} [\|\mathbf{e}_n^0 - f_{\theta}(\mathbf{e}_n^t, \mathbf{c}, t)\|_2^2]. \tag{7}$$

When an incredible item embedding $\mathbf{e}_j$ ($j \in \mathcal{I}_{unc}$) is used as the diffusion target (i.e., incredible diffusion target) during training, the diffusion loss encourages the prediction direction of the diffusion network f_θ to move closer to $\mathbf{e}_j$. Specifically, the MSE distance between the diffusion target and the output of f_θ will be smaller, indicating higher similarity.

In the inference stage, the generation process of diffusion recommenders can be expressed as:

$$\mathbf{e}_n^{t-1} = w_1 f_\theta(\mathbf{e}_n^t, \mathbf{c}, t) + w_2 \mathbf{e}_n^t + w_3 \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (8)$$

where $w_1 = \frac{\sqrt{\alpha_t} \beta_t}{1 - \alpha_t}$, $w_2 = \frac{\sqrt{\alpha_t}(1 - \alpha_t)}{1 - \alpha_t}$, and $w_3 = \sqrt{\frac{1 - \alpha_t}{1 - \alpha_t}}(1 - \alpha_t)$. This generation process is performed step by step, and the final embedding $\mathbf{e}_n^0$ is taken as the generation result, which then serves as the reference for item prediction and recommendations.

Let $\mathbf{e}_n^{t-1}$ denote the generated embedding at step $t - 1$ without using incredible diffusion target $\mathbf{e}_j$ during training. In such case, the parameters of the diffusion network are denoted as θ . Similarly, let $\hat{\mathbf{e}}_n^{t-1}$ denote the generated embedding at step $t - 1$ with $\mathbf{e}_j$ as the incredible diffusion target during training. In this case, the diffusion parameters are denoted as $\hat{\theta}$. We then calculate the difference in similarity between the normalized $\mathbf{e}_j$ and the normalized generated embeddings $\hat{\mathbf{e}}_n^{t-1}$ and $\mathbf{e}_n^{t-1}$ at step $t - 1$ as follows:

$$\begin{aligned} \Delta^{t-1} &= \text{sim}(\mathbf{e}_j, \hat{\mathbf{e}}_n^{t-1}) - \text{sim}(\mathbf{e}_j, \mathbf{e}_n^{t-1}) \\ &= [w_1 (f_{\hat{\theta}}(\hat{\mathbf{e}}_n^t, \mathbf{c}, t) - f_\theta(\mathbf{e}_n^t, \mathbf{c}, t)) + w_2 (\hat{\mathbf{e}}_n^t - \mathbf{e}_n^t) + w_3 (\boldsymbol{\epsilon}^t - \boldsymbol{\epsilon}^t)] \cdot \mathbf{e}_j^\top \\ &= w_1 (f_{\hat{\theta}}(\hat{\mathbf{e}}_n^t, \mathbf{c}, t) - f_\theta(\mathbf{e}_n^t, \mathbf{c}, t)) \cdot \mathbf{e}_j^\top + w_2 (\hat{\mathbf{e}}_n^t - \mathbf{e}_n^t) \cdot \mathbf{e}_j^\top. \end{aligned} \quad (9)$$

Here, we utilize the dot product to calculate the similarity. To avoid the interference from sampled noise, we use $\boldsymbol{\epsilon}^t$ to denote the sample noise $\boldsymbol{\epsilon}$ in step $t - 1$, and use it for both generation processes to control this variable.

When $t = T$, we have:

$$\begin{aligned} \Delta^{T-1} &= \text{sim}(\mathbf{e}_j, \hat{\mathbf{e}}_n^{T-1}) - \text{sim}(\mathbf{e}_j, \mathbf{e}_n^{T-1}) \\ &= w_1 (f_{\hat{\theta}}(\hat{\mathbf{e}}_n^T, \mathbf{c}, T) - f_\theta(\mathbf{e}_n^T, \mathbf{c}, T)) \cdot \mathbf{e}_j^\top + w_2 (\hat{\mathbf{e}}_n^T - \mathbf{e}_n^T) \cdot \mathbf{e}_j^\top \\ &= w_1 (f_{\hat{\theta}}(\boldsymbol{\epsilon}^T, \mathbf{c}, T) - f_\theta(\boldsymbol{\epsilon}^T, \mathbf{c}, T)) \cdot \mathbf{e}_j^\top + w_2 (\boldsymbol{\epsilon}^T - \boldsymbol{\epsilon}^T) \cdot \mathbf{e}_j^\top \\ &= w_1 (f_{\hat{\theta}}(\boldsymbol{\epsilon}^T, \mathbf{c}, T) \cdot \mathbf{e}_j^\top - f_\theta(\boldsymbol{\epsilon}^T, \mathbf{c}, T) \cdot \mathbf{e}_j^\top) \\ &= \underbrace{w_1}_{>0} \underbrace{(\text{sim}(\mathbf{e}_j, f_{\hat{\theta}}(\boldsymbol{\epsilon}^T, \mathbf{c}, T)) - \text{sim}(\mathbf{e}_j, f_\theta(\boldsymbol{\epsilon}^T, \mathbf{c}, T)))}_{>0} > 0. \end{aligned} \quad (10)$$

We control the process of two generations start from the same point $\hat{\mathbf{e}}_n^T = \mathbf{e}_n^T = \boldsymbol{\epsilon}^T$ for fair comparison. As mentioned earlier, the prediction direction of $f_{\hat{\theta}}$ is closer to $\mathbf{e}_j$ than that of f_θ . Hence, the MSE distance between the output of $f_{\hat{\theta}}$ and to $\mathbf{e}_j$ is smaller than that between the output of f_θ and $\mathbf{e}_j$. When the embeddings are normalized, a smaller MSE distance corresponds to a higher dot product similarity. Consequently, $\text{sim}(\mathbf{e}_j, f_{\hat{\theta}}(\boldsymbol{\epsilon}^T, \mathbf{c}, T)) - \text{sim}(\mathbf{e}_j, f_\theta(\boldsymbol{\epsilon}^T, \mathbf{c}, T)) > 0$. At the same time, $w_1 > 0$, therefore $\Delta^{T-1} > 0$. This indicates that, when starting from the same initial point, the generation result at step $T - 1$ produced by model $f_{\hat{\theta}}$, which has been trained with an incredible diffusion target, will be more similar to this incredible diffusion target.

When $t = T - 1$, we have:

$$\begin{aligned} \Delta^{T-2} &= w_1 (f_{\hat{\theta}}(\hat{\mathbf{e}}_n^{T-1}, \mathbf{c}, T - 1) - f_\theta(\mathbf{e}_n^{T-1}, \mathbf{c}, T - 1)) \cdot \mathbf{e}_j^\top + w_2 (\hat{\mathbf{e}}_n^{T-1} - \mathbf{e}_n^{T-1}) \cdot \mathbf{e}_j^\top \\ &= w_1 C_{T-1}^+ + w_2 (\text{sim}(\mathbf{e}_j, \hat{\mathbf{e}}_n^{T-1}) - \text{sim}(\mathbf{e}_j, \mathbf{e}_n^{T-1})) \\ &= w_1 C_{T-1}^+ + w_2 \Delta^{T-1}. \end{aligned} \quad (11)$$

As mentioned before, when a incredible item embedding $\mathbf{e}_j$ is taken for training, the diffusion loss will encourage the prediction direction of $f_{\hat{\theta}}$ closer to $\mathbf{e}_j$. At the same time, $\hat{\mathbf{e}}_n^{T-1}$ is closer to $\mathbf{e}_j$, as compared to that of $\mathbf{e}_n^{T-1}$. This further enforces $f_{\hat{\theta}}(\hat{\mathbf{e}}_n^{T-1}, \mathbf{c}, T - 1)$ more similar to $\mathbf{e}_j$, than that of $f_\theta(\mathbf{e}_n^{T-1}, \mathbf{c}, T - 1)$. Hence, the first term is a positive constant, and we denote it by C_{T-1}^+ .

Similarly, when $t < T - 1$, we have:

$$\begin{aligned} \Delta^{T-3} &= w_1 C_{T-2}^+ + w_2 \Delta^{T-2} \\ &= w_1 C_{T-2}^+ + w_2 (w_1 C_{T-1}^+ + w_2 \Delta^{T-1}) \\ &= w_1 C_{T-2}^+ + w_1 w_2 C_{T-1}^+ + w_2^2 \Delta^{T-1}. \end{aligned} \quad (12)$$

$$\begin{aligned}
\Delta^{T-4} &= w_1 C_{T-3}^+ + w_2 \Delta^{T-3} \\
&= w_1 C_{T-3}^+ + w_2 (w_1 C_{T-2}^+ + w_1 w_2 C_{T-1}^+ + w_2^2 \Delta^{T-1}) \\
&= w_1 C_{T-3}^+ + w_1 w_2 C_{T-2}^+ + w_1 w_2^2 C_{T-1}^+ + w_2^3 \Delta^{T-1}.
\end{aligned} \tag{13}$$

$$\begin{aligned}
\Delta^0 &= w_1 C_1^+ + w_1 w_2 C_2^+ + \dots + w_1 w_2^{T-2} C_{T-1}^+ + w_2^{T-1} \Delta^{T-1} \\
&= \underbrace{\sum_{m=1}^{T-1} w_1 w_2^{m-1} C_m^+}_{>0} + \underbrace{w_2^{T-1} \Delta^{T-1}}_{>0} \\
&= \text{sim}(\mathbf{e}_j, \hat{\mathbf{e}}_n^0) - \text{sim}(\mathbf{e}_j, \hat{\mathbf{e}}_n^0) \\
&> 0.
\end{aligned} \tag{14}$$

According to above analysis, the final generated result $\hat{\mathbf{e}}_n^0$ using diffusion network $f_{\hat{\theta}}$ is more similar with uncredible item embedding $\mathbf{e}_j$, as compared to the final generated result $\mathbf{e}_n^0$ using diffusion network f_{θ} . This indicates that when uncredible items are used as the diffusion targets during training, the model tends to generate outputs that carry more uncredible features, i.e., embeddings that are more similar to uncredible items.

Algorithm 1 Training of Disco

```

1: Input: Training dataset  $\mathcal{S}_{train} = \{(\mathbf{e}_n, \mathbf{e}_{neg}, s^{pre}, s^{unc})\}_{s=1}^{|\mathcal{S}_{train}|}$ , available uncredible item
   set  $\mathcal{I}_{unc}$ , trainable parameters  $\Theta$ , total diffusion steps  $T$ , learning rate  $\eta$ , variance schedules
    $\{\alpha_t\}_{t=0}^T$ .
2: Output: Optimized parameters  $\Theta$ .
3:  $\mathbf{F} = \text{Stack}(\{\mathbf{e}_i^{unc}\}_{i \in \mathcal{I}_{unc}})$  ▷ Construct uncredible feature matrix
4: repeat
5:    $(\mathbf{e}_n, \mathbf{e}_{neg}, s^{pre}, s^{unc}) \sim \mathcal{S}_{train}$  ▷ Sample training data
6:    $\mathbf{c}^{pre} = \text{Transformer}(s^{pre})$  ▷ Obtain preference-related condition
7:    $\mathbf{c}^{unc} = \text{Mean}(s^{unc})$  ▷ Obtain uncredible content-related condition
8:   Update F by Algorithm 3 ▷ Progressive uncredible feature matrix enhancement
9:    $[\mathbf{U}_1; \mathbf{U}_2], \mathbf{\Lambda}, \mathbf{V} = \text{SVD}(\mathbf{F}^\top)$  ▷ Construct null space of uncredible feature matrix
10:   $\tilde{\mathbf{e}}_n = \mathbf{e}_n \mathbf{U}_2 \mathbf{U}_2^\top$  ▷ Credible subspace projection for diffusion target  $\mathbf{e}_n$ 
11:   $\tilde{\mathbf{e}}_n = (\tilde{\mathbf{e}}_n + \mathbf{e}_n)/2$  ▷ Residual connection
12:   $t \sim \text{Uniform}(1, T)$  ▷ Sample diffusion step
13:   $\tilde{\mathbf{e}}_n^t = \sqrt{\alpha_t} \tilde{\mathbf{e}}_n^0 + \sqrt{1 - \alpha_t} \epsilon$  ▷ Add noise to the embedding of diffusion target
14:   $\mathbf{e}_{neg}^t = \sqrt{\alpha_t} \mathbf{e}_{neg}^0 + \sqrt{1 - \alpha_t} \epsilon$  ▷ Add noise to the embedding of negative preference item
15:   $\Theta \leftarrow \Theta - \eta \nabla_{\Theta} \mathcal{L}_{\text{Disco}}(\tilde{\mathbf{e}}_n, \mathbf{e}_{neg}, \mathbf{c}^{pre}, \mathbf{c}^{unc}, t, \Theta)$  ▷ Update parameters
16: until convergence
17: return  $\Theta$ 

```

Algorithm 2 Inference of Disco

```

1: Input: Test dataset  $\mathcal{S}_{test} = \{s^{pre}\}_{s=1}^{|\mathcal{S}_{test}|}$ , trained diffusion network parameters  $\theta \in \Theta$ , total
   reverse steps  $T$ , DDIM steps  $T'$ , variance schedules  $\{\alpha_t\}_{t=0}^T$ .
2: Output: A recommendation list for each user/sequence.
3:  $s^{pre} \sim \mathcal{S}_{test}$  ▷ Sample test sequence
4:  $\mathbf{c}^{pre} = \text{Transformer}(s^{pre})$  ▷ Obtain preference-related condition
5: for  $t' = T', \dots, 1$  do
6:    $t = \lfloor t' \times (T/T') \rfloor$  ▷ Calculate DDIM denoising step
7:    $\mathbf{e}_n^T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  ▷ Start from Gaussian noise
8:    $\mathbf{e}_n^{t-1} = \frac{\sqrt{\alpha_{t-1}(1-\alpha_t)}}{1-\alpha_t} f_{\theta}(\mathbf{e}_n^t, \mathbf{c}^{pre}, t) + \frac{\sqrt{\alpha_t(1-\alpha_{t-1})}}{1-\alpha_t} \mathbf{e}_n^t + \sqrt{\frac{1-\alpha_{t-1}}{1-\alpha_t}} (1-\alpha_t) \epsilon$  ▷
   Step-by-step generation
9: end for
10:  $\hat{y}_i = \mathbf{e}_n^0 \cdot \mathbf{e}_i^\top$  ▷ Calculate the matching score between a user/sequence and a candidate item  $\mathbf{e}_i$ 
11:  $\mathcal{R} = \{i | \text{TopK}(\hat{y}_i), i \in \mathcal{I}\}$  ▷ Select top K items with highest matching scores
12: return  $\mathcal{R}$ 

```

Algorithm 3 Progressive enhancement of incredible feature matrix

- 1: **Input:** Original incredible feature matrix $\mathbf{F}$, available incredible item set $\mathcal{I}_{unc}$, current iteration j , maximum selection ratio γ , maximum iteration m to reach γ .
 - 2: **Output:** Updated incredible feature matrix $\mathbf{F}$.
 - 3: $\text{UD}(i) = \frac{1}{|\mathcal{I}_{unc}|} \sum_{i' \in \mathcal{I}_{unc}} \cos(\mathbf{e}_i^{unc}, \mathbf{e}_{i'}^{unc})$ $\triangleright$ Calculate incredible degree of items in $\mathcal{I} \setminus \mathcal{I}_{unc}$
 - 4: $\text{ratio}(j) = \min(\gamma, \frac{j}{m}\gamma)$ $\triangleright$ Calculate the selection ratio at current iteration
 - 5: select $\lfloor |\mathcal{I} \setminus \mathcal{I}_{unc}| \cdot \text{ratio}(j) \rfloor$ items with highest incredible degree $\triangleright$ Select potential incredible items
 - 6: Add potential incredible items into $\mathcal{I}_{unc}$ $\triangleright$ Extension of incredible item set
 - 7: $\mathbf{F} = \text{Stack}(\{\mathbf{e}_i^{unc}\}_{i \in \mathcal{I}_{unc}})$ $\triangleright$ Enhancement of incredible feature matrix
 - 8: **return** $\mathbf{F}$
-

D DERIVATION OF EQUATION ??

In this section, we provide the derivation of Equation ?? . For simplicity, we only need to derive the first term, since the derivation of the second term follows the same procedure. The detailed derivation is as follows:

$$\begin{aligned}
-\mathbb{E}_q \left[\log \frac{p_\theta(\mathbf{e}_n^{0:T} | \mathbf{c}^{pre})}{q(\mathbf{e}_n^{1:T} | \mathbf{e}_n^0)} \right] &\stackrel{\textcircled{1}}{=} -\mathbb{E}_q \left[\log \frac{p(\mathbf{e}_n^T | \mathbf{c}^{pre}) p_\theta(\mathbf{e}_n^0 | \mathbf{e}_n^1, \mathbf{c}^{pre}) \prod_{t>1}^T p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre})}{q(\mathbf{e}_n^1 | \mathbf{e}_n^0) \prod_{t>1}^T q(\mathbf{e}_n^t | \mathbf{e}_n^{t-1}, \mathbf{e}_n^0)} \right] \\
&\stackrel{\textcircled{2}}{=} -\mathbb{E}_q \left[\log \frac{p(\mathbf{e}_n^T | \mathbf{c}^{pre}) p_\theta(\mathbf{e}_n^0 | \mathbf{e}_n^1, \mathbf{c}^{pre}) \prod_{t>1}^T p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre})}{q(\mathbf{e}_n^1 | \mathbf{e}_n^0) \prod_{t>1}^T \frac{q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0) q(\mathbf{e}_n^t | \mathbf{e}_n^0)}{q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^0)}} \right] \\
&\stackrel{\textcircled{3}}{=} -\mathbb{E}_q \left[\log \frac{p(\mathbf{e}_n^T | \mathbf{c}^{pre}) p_\theta(\mathbf{e}_n^0 | \mathbf{e}_n^1, \mathbf{c}^{pre}) \prod_{t>1}^T p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre})}{\frac{q(\mathbf{e}_n^1 | \mathbf{e}_n^0)}{q(\mathbf{e}_n^1 | \mathbf{e}_n^0)} \frac{q(\mathbf{e}_n^2 | \mathbf{e}_n^1)}{q(\mathbf{e}_n^2 | \mathbf{e}_n^1)} \dots \frac{q(\mathbf{e}_n^T | \mathbf{e}_n^{T-1})}{q(\mathbf{e}_n^T | \mathbf{e}_n^{T-1})} \prod_{t>1}^T q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0)} \right] \\
&\stackrel{\textcircled{4}}{=} -\mathbb{E}_q \left[\log p(\mathbf{e}_n^0 | \mathbf{e}_n^1, \mathbf{c}^{pre}) + \log \frac{p_\theta(\mathbf{e}_n^T)}{q(\mathbf{e}_n^T | \mathbf{e}_n^0)} + \log \frac{\prod_{t>1}^T p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre})}{\prod_{t>1}^T q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0)} \right] \\
&\stackrel{\textcircled{5}}{=} -\mathbb{E}_q [\log p(\mathbf{e}_n^0 | \mathbf{e}_n^1, \mathbf{c}^{pre})] - \mathbb{E}_q \left[\log \frac{p_\theta(\mathbf{e}_n^T)}{q(\mathbf{e}_n^T | \mathbf{e}_n^0)} \right] \\
&\quad - \sum_{t>1}^T \mathbb{E}_q \left[\log \frac{p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre})}{q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0)} \right] \\
&\stackrel{\textcircled{6}}{=} -\underbrace{\mathbb{E}_q [\log p(\mathbf{e}_n^0 | \mathbf{e}_n^1, \mathbf{c}^{pre})]}_{\text{reconstruction term}} + \underbrace{D_{\text{KL}}(q(\mathbf{e}_n^T | \mathbf{e}_n^0) \| p_\theta(\mathbf{e}_n^T))}_{\text{prior matching term}} \\
&\quad + \underbrace{\sum_{t>1}^T \mathbb{E}_q [D_{\text{KL}}(q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0) \| p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre}))]}_{\text{denoising matching term}}.
\end{aligned} \tag{15}$$

Equation $\textcircled{2}$ is derived through Bayes rule: $q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0) = \frac{q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0) q(\mathbf{e}_n^t | \mathbf{e}_n^0)}{q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^0)}$. Equation $\textcircled{4}$ is obtained since $p(\mathbf{e}_n^T | \mathbf{c}^{pre}) = p(\mathbf{e}_n^T)$ given $\mathbf{e}_n^T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, which is independent with condition $\mathbf{c}^{pre}$.

DMs generally optimize the denoising matching term $D_{\text{KL}}(q(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{e}_n^0) \| p_\theta(\mathbf{e}_n^{t-1} | \mathbf{e}_n^t, \mathbf{c}^{pre}))$ instead of the whole variational bound. Then, this denoising matching term can be derived into the optimization loss $\mathcal{L} = \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\frac{1}{2\sigma_t^2} \|\boldsymbol{\mu}_q(\mathbf{e}_n^t, \mathbf{e}_n^0) - \boldsymbol{\mu}_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t)\|_2^2 \right]$, by adding the condition $\mathbf{c}^{pre}$ into $\boldsymbol{\mu}_\theta(\mathbf{e}_n^t, t)$ in (Ho et al., 2020). Similar with (Pathak et al., 2023), $\boldsymbol{\mu}_q(\mathbf{e}_n^t, \mathbf{e}_n^0)$ is defined as

(Pathak et al., 2023):

$$\boldsymbol{\mu}_q(\mathbf{e}_n^t, \mathbf{e}_n^0) = \frac{\sqrt{\alpha_t}(1 - \bar{\alpha}_{t-1})\mathbf{e}_n^t + \sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)\mathbf{e}_n^0}{1 - \bar{\alpha}_t}. \quad (16)$$

In our model, $\boldsymbol{\mu}_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t)$ is defined as:

$$\boldsymbol{\mu}_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t) = \frac{\sqrt{\alpha_t}(1 - \bar{\alpha}_{t-1})\mathbf{e}_n^t + \sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t)}{1 - \bar{\alpha}_t}, \quad (17)$$

where $f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t)$ is the predicted $\mathbf{e}_n^0$ using the diffusion network f_θ .

Then, the optimization term can be rewritten as:

$$\begin{aligned} \mathcal{L} &= \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\frac{1}{2\sigma_t^2} \left\| \boldsymbol{\mu}_q(\mathbf{e}_n^t, \mathbf{e}_n^0) - \boldsymbol{\mu}_\theta(\mathbf{e}_n^t, t) \right\|_2^2 \right] \\ &= \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\frac{1}{2\sigma_t^2} \left\| \frac{\sqrt{\alpha_t}(1 - \bar{\alpha}_{t-1})\mathbf{e}_n^t + \sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)\mathbf{e}_n^0}{1 - \bar{\alpha}_t} \right. \right. \\ &\quad \left. \left. - \frac{\sqrt{\alpha_t}(1 - \bar{\alpha}_{t-1})\mathbf{e}_n^t + \sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t)}{1 - \bar{\alpha}_t} \right\|_2^2 \right] \\ &= \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\frac{1}{2\sigma_q^2(t)} \left\| \frac{\sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)}{1 - \bar{\alpha}_t} \mathbf{e}_n^0 - \frac{\sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)}{1 - \bar{\alpha}_t} f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t) \right\|_2^2 \right] \\ &= \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\frac{1}{2\sigma_q^2(t)} \left(\frac{\sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)}{1 - \bar{\alpha}_t} \right)^2 \left\| \mathbf{e}_n^0 - f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t) \right\|_2^2 \right]. \end{aligned} \quad (18)$$

In practice, the coefficient $\frac{1}{2\sigma_q^2(t)} \left(\frac{\sqrt{\bar{\alpha}_{t-1}}(1 - \alpha_t)}{1 - \bar{\alpha}_t} \right)^2$ is generally omitted (Ho et al., 2020). Hence, the optimization loss of our preference-related condition guided generation can be rewritten as $\mathcal{L}_{pre} = \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\left\| \mathbf{e}_n^0 - f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t) \right\|_2^2 \right]$. Similarly, the optimization loss of uncredible content-related condition guided generation is: $\mathcal{L}_{unc} = \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{unc}, t} \left[\left\| \mathbf{e}_n^0 - f_\theta(\mathbf{e}_n^t, \mathbf{c}^{unc}, t) \right\|_2^2 \right]$. Our `Disco` model aims to encourage the generation guided by preference-related condition and discourage the generation guided by uncredible content-related condition. To achieve this goal, the optimization objective is formulated as shown in Equation ??:

$$\mathcal{L} = \mathcal{L}_{pre} - \mathcal{L}_{unc} = \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{pre}, t} \left[\left\| \mathbf{e}_n^0 - f_\theta(\mathbf{e}_n^t, \mathbf{c}^{pre}, t) \right\|_2^2 \right] - \mathbb{E}_{\mathbf{e}_n^0, \mathbf{c}^{unc}, t} \left[\left\| \mathbf{e}_n^0 - f_\theta(\mathbf{e}_n^t, \mathbf{c}^{unc}, t) \right\|_2^2 \right]. \quad (19)$$

E THEORETICAL JUSTIFICATION OF CREDIBLE SUBSPACE PROJECTION

Our credible subspace projection is constructed using SVD. Applying SVD to the uncredible feature matrix $\mathbf{F}^\top$ yields the eigenvector matrix $\mathbf{U}$ and the diagonal eigenvalue matrix $\mathbf{\Lambda}$. $\mathbf{U}$ and $\mathbf{\Lambda}$ can be expressed as $\mathbf{U} = [\mathbf{U}_1; \mathbf{U}_2]$ and $\mathbf{\Lambda} = \begin{bmatrix} \mathbf{\Lambda}_1 & 0 \\ 0 & \mathbf{\Lambda}_2 \end{bmatrix}$. Correspondingly, $\mathbf{V}$ can be expressed as $\mathbf{V} = [\mathbf{V}_1; \mathbf{V}_2]$. All zero or near-zero singular values are contained in $\mathbf{\Lambda}_2$, and the corresponding eigenvectors are given by $\mathbf{U}_2$ and $\mathbf{V}_2$.

According to the principles of SVD, the following equation holds:

$$\mathbf{U}_2^\top \mathbf{F}^\top = \mathbf{U}_2^\top \mathbf{U}_1 \mathbf{\Lambda}_1 \mathbf{V}_1^\top. \quad (20)$$

Since the matrix $\mathbf{U}$ obtained from the SVD is an orthogonal matrix, we have:

$$\mathbf{U}_2^\top \mathbf{F}^\top = \underbrace{\mathbf{U}_2^\top \mathbf{U}_1}_{=0} \mathbf{\Lambda}_1 \mathbf{V}_1^\top = \mathbf{0}. \quad (21)$$

Our credible diffusion target is derived through $\tilde{\mathbf{e}}_n = \mathbf{e}_n \mathbf{U}_2 \mathbf{U}_2^\top$. Hence, we have:

$$\tilde{\mathbf{e}}_n \mathbf{F}^\top = \mathbf{e}_n \mathbf{U}_2 \underbrace{\mathbf{U}_2^\top \mathbf{F}^\top}_{=0} = \mathbf{0}. \quad (22)$$

This indicates that the derived credible diffusion target $\tilde{\mathbf{e}}_n$ is orthogonal to the uncredible feature matrix $\mathbf{F}$. In summary, our proposed credible subspace projection effectively removes uncredible content from the diffusion targets by ensuring that the projected targets lie orthogonally to the uncredible content.

F CASE STUDY

In this section, we conduct a case study using the GossipCop dataset to evaluate the effectiveness of `Disco`. The GossipCop dataset contains users’ interaction sequences with news articles, including both true news (i.e., credible items) and fake news (i.e., uncredible items). Specifically, we present the historical interaction sequences and recommendation lists for five users. The credible content items are marked in green, while uncredible items are marked in red. In addition, to illustrate the semantic relevance between content items, we utilize the same background color to highlight content with similar or related topics. From Table 9, we have the following observations:

- `Disco` demonstrates strong capability in delivering credible recommendations. Specifically, although all these users have interacted with uncredible items in their historical interaction sequences, the recommendation lists generated by `Disco` contain no uncredible content.
- `Disco` is capable of mitigating uncredible content while still preserving high recommendation accuracy. This is achieved by removing uncredible features while retaining users’ genuine preference-related information. For example, taking User4 as an example, this user had historically interacted with some news (including fake news) about the death of celebrities (highlighted in yellow). `Disco` can effectively capture this user’s genuine preference and recommend some content also in such topics. It is worth noting that User4 had interacted with fake news about the death of “Tom Petty”, and `Disco` recommends this user with a credible news article about the same event. This plays an important role in countering misinformation, as it helps users correct false impressions formed through prior exposure to uncredible content.

G USAGE CLAIM OF LARGE LANGUAGE MODELS

We only utilize ChatGPT for polishing the academic writing, with the prompt “Proofread the grammar and polish the writing of the given sentences”.

Table 9: Five cases showcasing the historical interaction sequences and the recommendation lists of five users sampled from GossipCop dataset. **Credible** refers to credible content (i.e., true news) and **Uncredible** refers to uncredible content (i.e., fake news). In a user sample, the texts marked by the same-color background refer to similar topics. “**ground truth**” means the corresponding recommended content items have been actually read by the user in the test set.

User1	Historical sequence	Credible : Justin Timberlake, Chris Stapleton release 'Say Something' song , video.	Uncredible : Nicole Kidman, Keith Urban: Secrets to a Successful Relationship.	Uncredible : Kendall Jenner Shades Scott Disick Over Photo With Sofia Richie and His Kids.	Uncredible : Grammy winners 2018: the complete list.	
	Recommendations	Credible 2018 Latin GRAMMY Awards Complete Winners List.	Credible : Weinstein Company Files for Bankruptcy and Revokes Nondisclosure Agreements.	Credible : Oscars: The Complete Winners List.	Credible : Pop superstar Lady Gaga has officially landed her first Las Vegas residency.	Credible (ground truth) : TV News Roundup: Netflix Reveals Fuller House Season 4 Premiere Date
User2	Historical sequence	Credible : 13 Nights Of Halloween 2017 Schedule: Full List of Movies.	Uncredible : Taylor Swift will reportedly keep her new album off streaming services like Spotify and Apple Music for a week.	Uncredible : Former NBC interviewer lashes out at Trump in an NYT op-ed for reportedly casting doubt on the authenticity of the infamous tape.	Credible : 'Big Little Lies' Season 2 News, Premiere Date & Cast.	
	Recommendations	Credible (ground truth) : Justin Timberlake Announces New Album Man of the Woods.	Credible : Seven-time and defending champion says she isn't quite ready to return after giving birth to daughter in September.	Credible : Pop superstar Lady Gaga has officially landed her first Las Vegas residency.	Credible : Jamie Lynn Spears' second child on the way will join big sister Maddie Briann.	Credible : "Good morning baby of mine, John Stamos' fiancé Caitlin McHugh wrote as she debuted her baby bump..."
User3	Historical sequence	Credible : Hugh Grant and Anna Eberstein's baby on the way joins their daughter .	Uncredible : The cancellation of the third Sex and the City film came with headline-making fallout something Sarah Jessica Parker struggled with	Uncredible : Selena Gomez has completed her treatment for depression and anxiety and is reported feeling	Credible : Congratulations are in order for Rachel McAdams the 39-year-old actress is reportedly going to be a first-time mom! Though she has not personally confirmed the baby news	
	Recommendations	Credible : All Chicago West Baby Photos Timeline.	Credible : Demi Lovato Says She Contemplated Suicide at Age 7.	Credible : 'Black Panther' is the most tweeted about movie ever.	Credible (ground truth) : His wife Faith Hill said the country star had been suffering from dehydration.	Credible : Tisha Campbell-Martin Files For Divorce From Husband of 21 Years
User4	Historical sequence	Uncredible : Caitlyn Jenner told Diane Sawyer that she had undergone the final surgery in her gender reassignment procedures on Friday night's 20/20 special.	Credible : Indiana police found the actress unresponsive after responding to a 911 call Saturday.	Credible : Roger Ailes, Former Fox News CEO, Dies At 77.	Uncredible : Tom Petty Dead : Celebrities React on Social Media Variety.	
	Recommendations	Credible : An emotional Celine Dion returned to the stage in Las Vegas on Tuesday night.	Credible (ground truth) : Rocker Tom Petty died Monday after being rushed to a Los Angeles hospital.	Credible : Hugh Hefner's death certificate from the Los Angeles County Department of Public Health.	Credible : The final season of Netflix's "House of Cards" keeps the secret of how Frank Underwood died until the very end.	Credible : Pauley Perrette announces she's leaving "NCIS" after 15 seasons.
User5	Historical sequence	Credible : Benjamin Glaze had never kissed a girl before Katy Perry tricked him during the ABC reboot of American Idol.	Uncredible : During her chat with Ryan Murphy Friday (March 16) for the opening night of PaleyFest in Los Angeles.	Credible : A longtime aerialist for the famed Cirque Du Soleil plummeted to his death in front of a horrified crowd in Florida on Saturday night while trying out a new act...	Uncredible : Justin Bieber 's struggling with his split from Selena Gomez as she's all smiles on her Australian vacation. Here's how the Biebs is coping with his...	
	Recommendations	Credible (ground truth) : Justin Bieber Wants to Be With Selena Gomez But Is Hanging With Baskin Champion.	Credible : The singer covered Ariana Grande's 'Just a Little Bit of Your Heart' in the arena where her concert was attacked...	Credible : Trevorrow helmed the rebooted franchise's first installment.	Credible : Voting closes at 5pm PT today (June 29) for this year's News' TV Scoop Awards...	Credible : Blake Shelton Gets His Palms Read With Jimmy Fallon, Jokes About Having Too Much Sex.

REFERENCES

- Zhuo Cai, Shoujin Wang, Victor W Chu, Usman Naseem, Yang Wang, and Fang Chen. Unleashing the potential of diffusion models towards diversified sequential recommendations. In *Proceedings of the 48th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 1476–1486, 2025.
- Jianxin Chang, Chen Gao, Yu Zheng, Yiqun Hui, Yanan Niu, Yang Song, Depeng Jin, and Yong Li. Sequential recommendation with graph neural networks. In *ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 378–387, 2021.
- Ziqiang Cui, Haolun Wu, Bowei He, Ji Cheng, and Chen Ma. Context matters: Enhancing sequential recommendation with context-aware diffusion-based contrastive learning. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management*, pp. 404–414, 2024a.
- Ziqiang Cui, Haolun Wu, Bowei He, Ji Cheng, and Chen Ma. Diffusion-based contrastive learning for sequential recommendation. *arXiv preprint arXiv:2405.09369*, 2024b.
- Yashar Deldjoo, Zhankui He, Julian McAuley, Anton Korikov, Scott Sanner, Arnau Ramisa, René Vidal, Maheswaran Sathiamoorthy, Atoosa Kasirzadeh, and Silvia Milano. A review of modern recommender systems using generative models (gen-recsys). In *Proceedings of the 30th ACM SIGKDD conference on Knowledge Discovery and Data Mining*, pp. 6448–6458, 2024.
- Miriam Fernandez, Alejandro Bellogín, and Iván Cantador. Analysing the effect of recommendation algorithms on the spread of misinformation. In *Proceedings of the 16th ACM Web Science Conference*, pp. 159–169, 2024.
- Jesse Harte, Wouter Zorgdrager, Panos Louridas, Asterios Katsifodimos, Dietmar Jannach, and Marios Fragkoulis. Leveraging large language models for sequential recommendation. In *Proceedings of the 17th ACM Conference on Recommender Systems*, pp. 1096–1102, 2023.
- Ruining He and Julian McAuley. Fusing similarity models with markov chains for sparse sequential recommendation. In *IEEE International Conference on Data Mining*, pp. 191–200, 2016.
- Balázs Hidasi, Alexandros Karatzoglou, Linas Baltrunas, and Domonkos Tikk. Session-based recommendations with recurrent neural networks. *arXiv preprint arXiv:1511.06939*, 2015.
- Jonathan Ho and Tim Salimans. Classifier-free diffusion guidance. In *NeurIPS 2021 Workshop on Deep Generative Models and Downstream Applications*.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. *Advances in neural information processing systems*, 33:6840–6851, 2020.
- Guoqing Hu, Zhengyi Yang, Zhibo Cai, An Zhang, and Xiang Wang. Generate and instantiate what you prefer: Text-guided diffusion for sequential recommendation. *arXiv preprint arXiv:2410.13428*, 2024.
- Wang-Cheng Kang and Julian McAuley. Self-attentive sequential recommendation. In *IEEE International Conference on Data Mining*, pp. 197–206, 2018.
- Jin Li, Shoujin Wang, Qi Zhang, Shui Yu, and Fang Chen. Generating with fairness: A modality-diffused counterfactual framework for incomplete multimodal recommendations. In *Proceedings of the ACM on Web Conference 2025*, pp. 2787–2798, 2025a.
- Wuchao Li, Rui Huang, Haijun Zhao, Chi Liu, Kai Zheng, Qi Liu, Na Mou, Guorui Zhou, Defu Lian, Yang Song, et al. Dimerec: a unified framework for enhanced sequential recommendation via generative diffusion models. In *Proceedings of the Eighteenth ACM International Conference on Web Search and Data Mining*, pp. 726–734, 2025b.
- Zihao Li, Aixin Sun, and Chenliang Li. Diffurec: A diffusion model for sequential recommendation. *ACM Transactions on Information Systems*, 42(3):1–28, 2023.

972 Jianghao Lin, Jiaqi Liu, Jiachen Zhu, Yunjia Xi, Chengkai Liu, Yangtian Zhang, Yong Yu, and
973 Weinan Zhang. A survey on diffusion models for recommender systems. *arXiv preprint*
974 *arXiv:2409.05033*, 2024.

975 Qidong Liu, Fan Yan, Xiangyu Zhao, Zhaocheng Du, Huifeng Guo, Ruiming Tang, and Feng Tian.
976 Diffusion augmentation for sequential recommendation. In *Proceedings of the 32nd ACM Inter-*
977 *national conference on information and knowledge management*, pp. 1576–1586, 2023.

979 Shuo Liu, An Zhang, Guoqing Hu, Hong Qian, and Tat-seng Chua. Preference diffusion for recom-
980 mendation. In *ICLR*, 2025a.

981 Ziwei Liu, Qidong Liu, Yejing Wang, Wanyu Wang, Pengyue Jia, Maolin Wang, Zitao Liu,
982 Yi Chang, and Xiangyu Zhao. Sigma: Selective gated mamba for sequential recommendation.
983 In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pp. 12264–12272,
984 2025b.

986 Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. *arXiv preprint*
987 *arXiv:1711.05101*, 2017.

988 Haokai Ma, Ruobing Xie, Lei Meng, Xin Chen, Xu Zhang, Leyu Lin, and Zhanhui Kang. Plug-
989 in diffusion model for sequential recommendation. In *Proceedings of the AAAI conference on*
990 *artificial intelligence*, volume 38, pp. 8886–8894, 2024a.

992 Haokai Ma, Ruobing Xie, Lei Meng, Yimeng Yang, Xingwu Sun, and Zhanhui Kang. Seedrec:
993 sememe-based diffusion for sequential recommendation. In *Proceedings of IJCAI*, pp. 1–9,
994 2024b.

995 Haokai Ma, Yimeng Yang, Lei Meng, Ruobing Xie, and Xiangxu Meng. Multimodal conditioned
996 diffusion model for recommendation. In *Companion Proceedings of the ACM on Web Conference*,
997 pp. 1733–1740, 2024c.

999 Zihan Ma, Minnan Luo, Yiran Hao, Zhi Zeng, Xiangzheng Kong, and Jiahao Wang. Bridging inter-
1000 ests and truth: Towards mitigating fake news with personalized and truthful recommendations. In
1001 *Proceedings of the 48th International ACM SIGIR Conference on Research and Development in*
1002 *Information Retrieval*, pp. 490–503, 2025.

1003 Royal Pathak, Francesca Spezzano, and Maria Soledad Pera. Understanding the contribution of rec-
1004 ommendation algorithms on misinformation recommendation and misinformation dissemination
1005 on social networks. *ACM Transactions on the Web*, 17(4):1–26, 2023.

1006 Yuanpeng Qu and Hajime Nobuhara. Intent-aware diffusion with contrastive learning for sequential
1007 recommendation. In *Proceedings of the 48th International ACM SIGIR Conference on Research*
1008 *and Development in Information Retrieval*, pp. 1552–1561, 2025.

1009 Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. Bpr: Bayesian
1010 personalized ranking from implicit feedback. In *Proceedings of the Twenty-Fifth Conference on*
1011 *Uncertainty in Artificial Intelligence*, pp. 452–461, 2009.

1012 Leheng Sheng, An Zhang, Yi Zhang, Yuxin Chen, Xiang Wang, and Tat-Seng Chua. Language
1013 representations can be what recommenders need: Findings and potentials. In *The Thirteenth*
1014 *International Conference on Learning Representations*, 2025.

1015 Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. In *Interna-*
1016 *tional Conference on Learning Representations*, 2021.

1017 Fei Sun, Jun Liu, Jian Wu, Changhua Pei, Xiao Lin, Wenwu Ou, and Peng Jiang. Bert4rec: Se-
1018 quential recommendation with bidirectional encoder representations from transformer. In *ACM*
1019 *International Conference on Information and Knowledge Management*, pp. 1441–1450, 2019.

1020 Jiayi Tang and Ke Wang. Personalized top-n sequential recommendation via convolutional sequence
1021 embedding. In *Proceedings of the eleventh ACM international conference on web search and data*
1022 *mining*, pp. 565–573, 2018.

1026 Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Niko-
1027 lay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open founda-
1028 tion and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.

1029

1030 Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez,
1031 Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural informa-*
1032 *tion processing systems*, 30, 2017.

1033

1034 Chenyang Wang, Weizhi Ma, Chong Chen, Min Zhang, Yiqun Liu, and Shaoping Ma. Sequential
1035 recommendation with multiple contrast signals. *ACM Transactions on Information Systems*, 41
1036 (1):1–27, 2023a.

1037

1038 Shoujin Wang, Liang Hu, Yan Wang, Longbing Cao, Quan Z. Sheng, and Mehmet Orgun. Sequential
1039 recommender systems: Challenges, progress and prospects. In *Twenty-Eighth International Joint*
1040 *Conference on Artificial Intelligence {IJCAI-19}*. International Joint Conferences on Artificial
1041 Intelligence Organization, 2019.

1042

1043 Shoujin Wang, Longbing Cao, Yan Wang, Quan Z Sheng, Mehmet A Orgun, and Defu Lian. A
1044 survey on session-based recommender systems. *ACM Computing Surveys*, 54(7):1–38, 2021.

1045

1046 Shoujin Wang, Xiaofei Xu, Xiuzhen Zhang, Yan Wang, and Wenzhuo Song. Veracity-aware and
1047 event-driven personalized news recommendation for fake news mitigation. In *Proceedings of the*
1048 *ACM web conference 2022*, pp. 3673–3684, 2022.

1049

1050 Shoujin Wang, Wentao Wang, Xiuzhen Zhang, Yan Wang, Huan Liu, and Fang Chen. A hierarchical
1051 and disentangling interest learning framework for unbiased and true news recommendation. In
1052 *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*,
1053 pp. 3200–3211, 2024a.

1054

1055 Wenjie Wang, Yiyang Xu, Fuli Feng, Xinyu Lin, Xiangnan He, and Tat-Seng Chua. Diffusion rec-
1056ommender model. In *ACM SIGIR Conference on Research and Development in Information*
1057 *Retrieval*, pp. 832–841, 2023b.

1058

1059 Yu Wang, Zhiwei Liu, Liangwei Yang, and Philip S Yu. Conditional denoising diffusion for sequen-
1060 tial recommendation. In *Pacific-Asia Conference on Knowledge Discovery and Data Mining*, pp.
1061 156–169, 2024b.

1062

1063 Ting-Ruen Wei and Yi Fang. Diffusion models in recommendation systems: A survey. *arXiv*
1064 *preprint arXiv:2501.10548*, 2025.

1065

1066 Zihao Wu, Xin Wang, Hong Chen, Kaidong Li, Yi Han, Lifeng Sun, and Wenwu Zhu. Diff4rec:
1067 Sequential recommendation with curriculum-scheduled diffusion augmentation. In *Proceedings*
1068 *of the 31st ACM international conference on multimedia*, pp. 9329–9335, 2023.

1069

1070 Wenjia Xie, Rui Zhou, Hao Wang, Tingjia Shen, and Enhong Chen. Bridging user dynamics: Trans-
1071 forming sequential recommendations with schrödinger bridge and diffusion models. In *Proceed-*
1072 *ings of the 33rd ACM International Conference on Information and Knowledge Management*, pp.
1073 2618–2628, 2024.

1074

1075 Xu Xie, Fei Sun, Zhaoyang Liu, Shiwen Wu, Jinyang Gao, Jiandong Zhang, Bolin Ding, and Bin
1076 Cui. Contrastive learning for sequential recommendation. In *IEEE International Conference on*
1077 *Data Mining*, pp. 1259–1273, 2022.

1078

1079 Zhengyi Yang, Xiangnan He, Jizhi Zhang, Jiancan Wu, Xin Xin, Jiawei Chen, and Xiang Wang. A
generic learning framework for sequential recommendation with distribution shifts. In *Proceed-*
ings of the 46th International ACM SIGIR Conference on Research and Development in Informa-
tion Retrieval, pp. 331–340, 2023a.

Zhengyi Yang, Jiancan Wu, Zhicai Wang, Xiang Wang, Yancheng Yuan, and Xiangnan He. Generate
what you prefer: Reshaping sequential recommendation via guided diffusion. In *Conference on*
Neural Information Processing Systems, volume 36, 2023b.

1080 Ghim-Eng Yap, Xiao-Li Li, and Philip S Yu. Effective next-items recommendation via personal-
1081 ized sequential pattern mining. In *International conference on database systems for advanced*
1082 *applications*, pp. 48–64. Springer, 2012.

1083 Zhenrui Yue, Yueqi Wang, Zhankui He, Huimin Zeng, Julian McAuley, and Dong Wang. Linear
1084 recurrent units for sequential recommendation. In *ACM International Conference on Web Search*
1085 *and Data Mining*, pp. 930–938, 2024.

1086 Jujia Zhao, Wang Wenjie, Yiyang Xu, Teng Sun, Fuli Feng, and Tat-Seng Chua. Denoising diffusion
1087 recommender model. In *ACM SIGIR Conference on Research and Development in Information*
1088 *Retrieval*, pp. 1370–1379, 2024.

1089
1090
1091
1092
1093
1094
1095
1096
1097
1098
1099
1100
1101
1102
1103
1104
1105
1106
1107
1108
1109
1110
1111
1112
1113
1114
1115
1116
1117
1118
1119
1120
1121
1122
1123
1124
1125
1126
1127
1128
1129
1130
1131
1132
1133