

Mitigating Consistency Bias in Recommender Systems via Transformer-Based Dynamic Fuzzy Bias Adjustment

Sun Wen

School of Computer and Artificial Intelligence, Nanjing University of Finance and Economics, Nanjing, China
sw2232696283@163.com

Abstract: Recommender systems provide personalized services by learning from users' historical interactions, but recorded feedback may combine intrinsic preference with external influence and persistent rating habits. This mixture can produce consistency bias, in which users remain close to historical item evaluations or their own previous scoring patterns even when current interests have changed. To address this problem, this paper proposes a consistency-bias mitigation framework that combines Transformer-based sequential modeling with dynamic fuzzy bias adjustment. A Transformer encoder first captures long-term dependencies and short-term changes in users' interaction sequences. The model then measures rating variation between adjacent temporal stages, adaptively refines temporal windows, and converts the variation signal into continuous fuzzy weights. The adjusted historical representation is fused with the sequential interest representation for candidate-item prediction. Experiments on MovieLens100K, Amazon Beauty, and Amazon Toys compare the framework with collaborative filtering, neural recommendation, recurrent sequential recommendation, and Transformer-based baselines. Macro-level comparisons, dataset-wise analysis, ablation tests, robustness experiments, parameter sensitivity, and a case study show that the proposed method improves Recall and NDCG while degrading more slowly when biased interactions are introduced. These results suggest that temporal change detection and adaptive weighting can reduce the excessive influence of repeated historical patterns without discarding useful behavioral evidence.

Keywords: Recommender System, Consistency Bias, Transformer, Dynamic Fuzzy Adjustment, Sequential Recommendation, Bias Mitigation

1. Introduction

With the rapid growth of e-commerce, streaming media, and online content services, users face an increasingly large information space. Recommender systems have therefore become essential for reducing information overload and supporting content discovery. Traditional methods, including collaborative filtering and matrix factorization, infer preferences from historical user-item interactions and often assume that observed behavior directly reflects intrinsic interest.

In practice, feedback may also be shaped by existing item ratings, similar users, popularity trends, and persistent personal scoring habits. These factors can create consistency bias, causing behavior to remain close to prior information or stable historical patterns. If every interaction is treated as equally reliable, external influence may become entangled with individual preference [1,2].

Consistency bias is related to, but distinct from, popularity bias. Popularity describes an item's overall attention, whereas consistency reflects how strongly an individual response follows prior information or group feedback. Because users may react differently to the same item, bias correction should be user- and time-sensitive rather than global.

Sequential recommendation provides a suitable basis for this task because interactions are ordered chronologically. GRU-based models summarize behavior through recurrent states, while Transformer-based methods capture long-range dependencies through self-attention [3,4,5]. However, most models do not assess whether earlier interactions remain reliable indicators of current interest. Stable behavior may represent genuine preference or repeated conformity, whereas rapid change may indicate that recent evidence deserves greater weight.

This paper combines Transformer-based sequence encoding with dynamic fuzzy bias adjustment. Rather than removing all external influence, the framework estimates changes in rating tendency and adaptively controls the contribution of historical information. The main contributions are:

- (1) A Transformer-based encoder captures short-term changes and long-term dependencies in user behavior.
- (2) A dynamic mechanism measures rating variation, adjusts temporal windows, and generates continuous fuzzy weights.
- (3) A bias-aware fusion strategy incorporates adjusted historical evidence while preserving personalized interest.
- (4) Experiments on three public datasets evaluate accuracy, ablation, robustness, sensitivity, and qualitative behavior.

Section 2 reviews related work, Section 3 presents the method, Section 4 describes the experiments, Section 5 discusses the results, and Section 6 concludes the paper.

2. Related Work

2.1. Sequential Recommendation

Early sequential recommenders modeled local transitions through Markov chains or recurrent networks. FPMC combined personalized transitions with long-term latent preference [6], and GRU4Rec introduced gated recurrent units for session-based recommendation [5]. Transformer models then became prominent because self-attention captures distant relations directly. SASRec applies unidirectional attention [3], BERT4Rec uses bidirectional context and masked-item prediction [4], and S3-Rec adds self-supervised objectives linking items, attributes, and sequences [7].

Although these models learn strong sequential representations, they usually treat all recorded interactions as homogeneous preference signals. Group influence or rating inertia may therefore be preserved and amplified.

2.2. Bias-Aware Recommendation

Recommendation bias includes exposure, position, popularity, selection, and conformity effects [1]. Reweighting methods adjust observed samples, while causal methods separate preference from confounding factors. DICE disentangles interest and conformity [2], MACR uses counterfactual reasoning to reduce popularity influence [8], and PDA selectively reintroduces useful popularity information after deconfounded training [9]. TIDE separates item quality from time-varying conformity [10], while session-level work distinguishes short-term interest from popular-choice conformity [11].

These studies show that external influence should not always be removed, because it may also reflect item quality or the user's decision environment. The key is to determine when such evidence should be retained or weakened.

2.3. Side Information and Research Gap

Side-information methods use categories, brands, text, and attributes to improve representation learning. NOVA modifies attention without replacing item representations [12], DIF-SR decouples item and attribute attention [13], and ASIF improves alignment between item identity and auxiliary features [14]. DCRRec combines sequential and collaborative signals through debiased contrastive learning [15].

Three gaps remain. Existing bias-aware methods often rely on static statistics, sequence models may interpret repeated behavior as stable interest, and fixed reweighting cannot adapt to different rates of preference change. The proposed framework addresses these limitations through long-range sequence modeling, user-specific change detection, and continuous fuzzy weighting.

3. Methodology

The proposed framework contains four components: sequence construction, Transformer-based interest encoding, dynamic fuzzy bias adjustment, and bias-aware fusion for prediction. The Transformer

representation remains the principal description of user preference, while the dynamically adjusted historical signal is introduced as an auxiliary residual. This design prevents the bias component from completely replacing the personalized sequence representation.

3.1. Problem Definition

Let U denote the user set and I denote the item set. The chronological interaction sequence of user u is defined as follows:

$$S_u = \{(i_1, r_1, t_1), (i_2, r_2, t_2), \dots, (i_n, r_n, t_n)\} \quad (1)$$

where i denotes an interacted item, r denotes an explicit rating or transformed feedback value, and t denotes the interaction time. Given the prefix of S_u , the model predicts either the preference score for a candidate item or the probability that the item will be selected next.

3.2. Transformer-Based Interest Representation

Each interaction is represented by the sum of item, rating, temporal-position, and user embeddings. Positional information preserves the chronological order that would otherwise be lost in self-attention. The embedded sequence is processed by stacked Transformer encoder layers:

$$H_u = \text{Transformer}(E(S_u)) \quad (2)$$

where $E(\cdot)$ is the embedding function and H_u contains contextualized hidden states. The final valid hidden state h_u is used as the current interest representation. Multi-head attention allows the model to focus on multiple behavioral relations, such as repeated item categories, recent preference shifts, and long-term consumption patterns.

Unlike recurrent networks, the Transformer does not compress the complete sequence into a single recurrent transition. It can directly compare distant interactions and is therefore suitable for identifying whether recent behavior is consistent with or different from earlier rating patterns.

3.3. Dynamic Fuzzy Bias Adjustment

The historical sequence is divided into temporal stages. The initial stage length is fixed, but adjacent stages can be split or merged according to rating variation. The normalized change rate between adjacent stages is defined as:

$$\Delta_u^k = |R_u^k - R_u^{k-1}| / (R_{\max} - R_{\min}) \quad (3)$$

where R_u^k is the average rating of user u in stage k . A large change rate indicates that the current stage differs from the preceding historical pattern. The model therefore uses a finer window to preserve recent evidence. When the change rate is small, adjacent windows are merged to reduce redundant segmentation and prevent repeated stable behavior from dominating the representation.

The change rate is converted into a continuous weight through a fuzzy Sigmoid mapping:

$$w_u^k = \sigma(\alpha \Delta_u^k + \beta) \quad (4)$$

where α controls sensitivity and β controls the baseline contribution of a stage. Compared with a hard threshold, the continuous mapping avoids abrupt changes in the contribution of neighboring temporal stages. The parameters are learned or selected on the validation set, allowing the model to adapt the adjustment strength to different datasets.

3.4. Bias-Aware Fusion and Prediction

For each temporal stage, the hidden states are aggregated into a stage representation g_u^k . The dynamically adjusted historical signal is obtained by weighting these stage representations:

$$b_u = (\sum_{k=1}^K w_u^k g_u^k) / (\sum_{k=1}^K w_u^k + \varepsilon) \quad (5)$$

The adjusted signal b_u is projected into the same space as the Transformer representation and introduced through residual fusion:

$$\tilde{h}_u = h_u + W_b b_u \quad (6)$$

where W_b is a trainable projection matrix. The candidate-item score is calculated by the inner product

between the bias-aware user representation and the candidate-item embedding:

$$s(u,i) = \tilde{\mathbf{h}}_u^T \mathbf{e}_i \quad (7)$$

For next-item recommendation, the model is optimized with full-item cross-entropy and L2 regularization. For explicit-rating prediction, a mean squared error term can be added as an auxiliary objective. The use of a residual connection ensures that adjusted historical evidence supplements rather than replaces the Transformer interest representation.

3.5. Complexity Discussion

For a maximum sequence length L and hidden dimension d , the main Transformer computation has time complexity $O(L^2d)$, while temporal-stage aggregation and fuzzy weighting require approximately $O(Ld)$. Therefore, the additional bias-adjustment module does not change the dominant asymptotic complexity of the backbone. The extra parameters mainly arise from the stage-projection and fusion layers, which are small compared with the item-embedding matrix.

4. Experimental Design

4.1. Datasets

Experiments are conducted on MovieLens100K, Amazon Beauty, and Amazon Toys. MovieLens100K contains explicit ratings and timestamps and supports direct analysis of rating tendency. Amazon Beauty and Amazon Toys contain sparse product interactions and richer domain variation. Their high sparsity makes them suitable for testing whether the proposed method can use temporal signals without relying on dense rating histories. Dataset statistics are shown in Table 1.

Table 1: Statistics of the datasets.

Dataset	Users	Items	Interactions	Sparsity
MovieLens100K	943	1,682	100,000	93.70%
Amazon Beauty	22,363	12,101	198,502	99.93%
Amazon Toys	19,412	11,924	167,597	99.93%

The datasets cover both explicit-rating recommendation and sparse implicit-feedback recommendation. This combination is useful because consistency bias may be expressed as repeated rating habits in explicit feedback and as repeated selection of popular or group-consistent items in implicit feedback.

4.2. Data Preprocessing

Interactions are sorted by timestamp for each user. User and item identifiers are mapped to consecutive indices, and sequences shorter than five interactions are removed. For MovieLens100K, ratings remain on the original 1-5 scale and are also normalized when calculating temporal variation. For the Amazon datasets, review interactions are treated as positive implicit feedback.

A leave-one-out protocol is used: the last interaction is used for testing, the second-to-last interaction is used for validation, and all earlier interactions form the training sequence. The maximum sequence length is set to 50. Short sequences are left-padded, while longer sequences retain the most recent interactions.

4.3. Baseline Models

The baselines cover static, neural, recurrent, and Transformer-based recommendation. MF represents conventional latent-factor recommendation. NeuMF models nonlinear user-item relations through neural networks [16]. GRU4Rec is a recurrent sequential recommender [5]. SASRec and BERT4Rec represent unidirectional and bidirectional Transformer-based sequence modeling [3,4]. These baselines allow the experiment to distinguish gains from temporal modeling from gains produced by the proposed bias-adjustment mechanism.

4.4. Parameter Settings

All models use the same data splits and evaluation candidates. Hyperparameters are selected according to validation NDCG@10. The principal settings of the proposed method are listed in Table 2.

Table 2: Main hyperparameter settings.

Parameter	Value	Parameter	Value
Embedding dimension	64	Maximum sequence length	50
Transformer layers	2	Attention heads	4
Batch size	256	Learning rate	0.001
Initial time window	30 days	Dropout rate	0.2
Sensitivity alpha	1.5	Base bias beta	-0.3
Optimizer	Adam	Early-stopping patience	10

The initial temporal window is used only as a starting point. The dynamic module may split or merge stages according to the observed change rate. Models are trained for at most 200 epochs, and the checkpoint with the best validation NDCG@10 is used for testing.

4.5. Evaluation Metrics and Research Questions

Ranking performance is evaluated with Recall@10, Recall@20, NDCG@10, and NDCG@20. Recall measures whether the ground-truth item is retrieved, whereas NDCG additionally rewards a higher ranking position. For MovieLens100K, MAE and RMSE are calculated as supplementary rating-prediction metrics.

The experiments address four research questions. RQ1 evaluates whether the complete method improves overall recommendation performance. RQ2 examines the contribution of each component. RQ3 tests robustness when the proportion of consistency-biased interactions increases. RQ4 analyzes whether the fuzzy-weight parameters behave stably within a reasonable range.

5. Results and Analysis

5.1. Overall Performance Comparison

Table 3 reports macro-averaged results across the three datasets. The average gives equal importance to each domain and avoids allowing the largest dataset to dominate the comparison.

Table 3: Overall performance comparison of different recommendation models.

Model	Recall@10	Recall@20	NDCG@10	NDCG@20
MF	0.0721	0.1054	0.0387	0.0479
NeuMF	0.0785	0.1128	0.0415	0.0512
GRU4Rec	0.0836	0.1187	0.0443	0.0538
SASRec	0.0916	0.1279	0.0486	0.0572
BERT4Rec	0.0942	0.1308	0.0501	0.0591
Proposed Method	0.1015	0.1396	0.0554	0.0643

Sequential models outperform MF and NeuMF, confirming the value of temporal order. The proposed method achieves the best performance on all four metrics. Relative to BERT4Rec, it improves Recall@10, Recall@20, NDCG@10, and NDCG@20 by 7.75%, 6.73%, 10.58%, and 8.80%, respectively. The larger improvement in NDCG suggests that the method not only retrieves more relevant items but also places them closer to the top of the list.

5.2. Dataset-Wise Performance

To examine whether the gains are concentrated in a single domain, Table 4 compares the strongest baseline with the proposed method on each dataset.

Table 4: Dataset-wise comparison with the strongest baseline.

Dataset	Method	Recall@10	NDCG@10	Relative NDCG Gain
MovieLens100K	BERT4Rec	0.1124	0.0652	-
MovieLens100K	Proposed	0.1218	0.0715	9.66%

Amazon Beauty	BERT4Rec	0.0735	0.0361	-
Amazon Beauty	Proposed	0.0796	0.0403	11.63%
Amazon Toys	BERT4Rec	0.0968	0.0490	-
Amazon Toys	Proposed	0.1031	0.0544	11.02%

The proposed method improves NDCG@10 on all three datasets. The gain on MovieLens100K shows that the framework can use explicit rating variation directly. The larger relative gains on the Amazon datasets indicate that temporal weighting remains useful in sparse implicit-feedback settings, where repeated selection of popular or group-consistent items can otherwise dominate a short user sequence.

5.3. Ablation Study

The Transformer encoder, dynamic temporal adjustment, and adaptive fuzzy weighting are removed separately. Table 5 reports the results.

Table 5: Ablation results of the proposed model.

Variant	Recall@10	Recall@20	NDCG@10	NDCG@20
w/o Transformer	0.0927	0.1285	0.0490	0.0575
w/o Dynamic Adjustment	0.0961	0.1321	0.0512	0.0604
w/o Adaptive Weight	0.0980	0.1357	0.0528	0.0621
Full Model	0.1015	0.1396	0.0554	0.0643

Removing the Transformer causes the largest decline because long-range dependency modeling is the foundation of the user representation. Removing dynamic temporal adjustment also produces a clear decrease, showing that fixed segmentation cannot adequately represent users whose behavior changes at different speeds. The smaller but consistent decline without adaptive weighting indicates that continuous fuzzy control is preferable to simply detecting a change and applying a fixed correction.

5.4. Robustness Analysis

Consistency-biased interactions are simulated in 10% to 30% of the test data. For explicit ratings, selected scores are shifted toward the historical average of the item. For implicit feedback, selected target interactions are replaced with items that were popular in the same period. Recall@10 under different bias ratios is shown in Table 6.

Table 6: Robustness under different proportions of consistency-biased interactions.

Bias Ratio	MF	SASRec	Proposed Method
0%	0.0721	0.0916	0.1015
10%	0.0684	0.0879	0.0992
20%	0.0635	0.0816	0.0958
30%	0.0572	0.0734	0.0911

All models decline as biased interactions increase, but the proposed method degrades more slowly. At a 30% bias ratio, MF and SASRec decrease by 20.67% and 19.87% relative to their unbiased results, whereas the proposed method decreases by 10.25%. This result supports the view that the method does not merely improve accuracy under a clean setting; it also limits the repeated amplification of biased historical patterns.

5.5. Parameter Sensitivity

The sensitivity coefficient alpha controls how strongly the change rate affects the fuzzy weight. Table 7 reports the macro-averaged results when alpha varies and beta remains fixed at -0.3.

Table 7: Sensitivity analysis of the fuzzy-weight coefficient.

Alpha	Recall@10	NDCG@10	Observation
0.5	0.0981	0.0528	Adjustment is too weak
1.0	0.1002	0.0545	Stable improvement
1.5	0.1015	0.0554	Best validation setting
2.0	0.1009	0.0549	Slight overreaction
3.0	0.0988	0.0536	Recent changes dominate

Performance is relatively stable between alpha=1.0 and alpha=2.0, indicating that the method does not depend on an extremely narrow parameter range. A very small value underuses temporal change,

while a very large value makes the representation overly sensitive to short-term fluctuations.

5.6. Case Study and Discussion

A user whose historical ratings were concentrated between 4 and 5 but whose recent ratings for science-fiction films declined was selected for qualitative analysis. A conventional sequential model continued to treat early high ratings as dominant evidence. The proposed method detected that the average rating decreased from 4.6 to 3.2, produced a change rate of 0.35, and reduced the contribution of repetitive early high-rating stages. The final recommendation list therefore contained fewer items closely matching the user's earlier science-fiction pattern and more drama and mystery items related to recent interactions.

The results support three observations. First, temporal modeling alone is not sufficient because a Transformer can still amplify repeated biased evidence. Second, dynamic adjustment should be continuous rather than binary, since user behavior usually changes gradually. Third, the usefulness of external influence is context dependent. Stable historical behavior should be retained when it remains predictive, but its contribution should be reduced when recent evidence indicates a meaningful change.

The proposed framework is intentionally lightweight and can be attached to other sequential backbones. However, the current definition of consistency bias is based mainly on temporal rating behavior. It does not fully distinguish item anchoring, group conformity, and genuine stable preference. This limitation motivates richer causal and semantic modeling in future work.

6. Conclusions and Future Work

This paper studies consistency bias in recommender systems and proposes a framework that combines Transformer-based sequence encoding with dynamic fuzzy bias adjustment. The method measures temporal variation in user feedback, adaptively refines temporal stages, and converts behavioral change into continuous weights. The resulting historical signal is fused with the Transformer representation through a residual mechanism, preserving personalized interest while limiting the dominance of repetitive historical patterns.

Experiments on MovieLens100K, Amazon Beauty, and Amazon Toys show that the proposed method achieves higher Recall and NDCG than collaborative filtering, neural recommendation, recurrent, and Transformer-based baselines. Dataset-wise results demonstrate consistent gains across explicit and implicit feedback. Ablation experiments verify the contribution of each module, robustness tests show a slower performance decline under simulated biased interactions, and sensitivity analysis indicates stable behavior within a practical parameter range.

Future work can improve the framework in three directions. First, review text, item attributes, social relations, and multimodal content can be used to explain the semantic source of consistency bias. Second, causal inference may help distinguish interest-driven interaction from group influence and exposure effects. Third, efficient attention and prototype-based temporal aggregation can reduce computation on larger datasets. Evaluating diversity, long-tail coverage, and user-level fairness would also provide a broader understanding of the practical impact of bias adjustment.

References

- [1] Chen J., Dong H., Wang X., et al. *Bias and Debias in Recommender System: A Survey and Future Directions*. *ACM Transactions on Information Systems*, 2023, 41(3): 1-39.
- [2] Zheng Y., Gao C., Li X., et al. *Disentangling User Interest and Conformity for Recommendation with Causal Embedding*. *Proceedings of WWW*, 2021: 2980-2991.
- [3] Kang W. C., McAuley J. *Self-Attentive Sequential Recommendation*. *Proceedings of ICDM*, 2018: 197-206.
- [4] Sun F., Liu J., Wu J., et al. *BERT4Rec: Sequential Recommendation with Bidirectional Encoder Representations from Transformer*. *Proceedings of CIKM*, 2019: 1441-1450.
- [5] Hidasi B., Karatzoglou A., Baltrunas L., Tikk D. *Session-based Recommendations with Recurrent Neural Networks*. *Proceedings of ICLR*, 2016.
- [6] Rendle S., Freudenthaler C., Schmidt-Thieme L. *Factorizing Personalized Markov Chains for Next-Basket Recommendation*. *Proceedings of WWW*, 2010: 811-820.
- [7] Zhou K., Wang H., Zhao W. X., et al. *S3-Rec: Self-Supervised Learning for Sequential*

- Recommendation with Mutual Information Maximization. Proceedings of CIKM, 2020: 1893-1902.*
- [8] Wei T., Feng F., Chen J., et al. *Model-Agnostic Counterfactual Reasoning for Eliminating Popularity Bias in Recommender System. Proceedings of KDD, 2021: 1791-1800.*
- [9] Zhang Y., Feng F., He X., et al. *Causal Intervention for Leveraging Popularity Bias in Recommendation. Proceedings of SIGIR, 2021: 11-20.*
- [10] Zhao Z., Chen J., Zhou S., et al. *Popularity Bias Is Not Always Evil: Disentangling Benign and Harmful Bias for Recommendation. IEEE Transactions on Knowledge and Data Engineering, 2023, 35(10): 9920-9931.*
- [11] Liu Q., Tian F., Zheng Q., Wang Q. *Disentangling Interest and Conformity for Eliminating Popularity Bias in Session-based Recommendation. Knowledge and Information Systems, 2023, 65(6): 2645-2664.*
- [12] Liu C., Li X., Cai G., et al. *Non-invasive Self-attention for Side Information Fusion in Sequential Recommendation. arXiv:2103.03578, 2021.*
- [13] Xie Y., Zhou P., Kim S. *Decoupled Side Information Fusion for Sequential Recommendation. Proceedings of SIGIR, 2022.*
- [14] Wang S., Shen B., Min X., et al. *Aligned Side Information Fusion Method for Sequential Recommendation. Companion Proceedings of the Web Conference, 2024: 112-120.*
- [15] Yang Y., Huang C., Xia L., et al. *Debiased Contrastive Learning for Sequential Recommendation. Proceedings of the Web Conference, 2023: 1063-1073.*
- [16] He X., Liao L., Zhang H., et al. *Neural Collaborative Filtering. Proceedings of WWW, 2017: 173-182.*