

Improving RecFormer with Sliding Window-Based Data Augmentation

Sungho Choi*, Minho Kim**, Namgyu Kim***

*Graduate Student, Graduate School of Business IT, Kookmin University, Seoul, Korea

**Presedent, ECNE Inc., Seoul, Korea

***Professor, Graduate School of Business IT, Kookmin University, Seoul, Korea

[Abstract]

With the rapid growth of the online content market, personalized recommendation systems have become increasingly important, yet Cold-start issues caused by data scarcity remain a critical challenge in new IP content environments. While existing ID-based sequential recommendation models fail to leverage textual attributes, RecFormer—built on the Longformer architecture—also suffers from performance degradation when training data is sparse. This paper proposes a Sliding Window-based data augmentation method for RecFormer, generating multiple training samples from a single user interaction sequence. Experiments on the Amazon Scientific Review dataset show that the proposed method consistently outperforms the baseline RecFormer across all metrics, with greater gains observed under more severe data sparsity conditions.

▶ **Key words:** RecFormer, Sliding Window, Data Augmentation, Cold-start, Text-based Recommendation

[요 약]

온라인 콘텐츠 시장의 성장과 함께 개인화 추천 시스템의 중요성이 높아지고 있으나, 신규 IP 콘텐츠 환경에서는 데이터 부족으로 인한 Cold-start 문제가 심화되고 있다. 기존 ID 기반 순차 추천 모델은 텍스트 속성을 활용하지 못하며, Longformer 기반의 텍스트 추천 모델인 RecFormer 역시 학습 데이터가 희소한 환경에서는 성능이 저하되는 한계가 있다. 이에 본 논문에서는 Sliding Window 기반 데이터 증강 기법을 RecFormer에 적용하여, 하나의 사용자 시퀀스로부터 다수의 학습 샘플을 생성하는 방법을 제안한다. Amazon Scientific Review 데이터셋을 활용한 실험 결과, 제안 방법은 모든 평가 지표에서 기존 RecFormer 대비 일관된 성능 향상을 보였으며, 데이터가 희소할수록 향상 폭이 더 크게 나타남을 확인하였다.

▶ **주제어:** RecFormer, 슬라이딩 윈도우, 데이터 증강, 콜드 스타트, 텍스트 기반 추천

-
- First Author: Sungho Choi, Corresponding Author: Minho Kim
 - *Sungho Choi (unlver@kookmin.ac.kr), Graduate School of Business IT, Kookmin University
 - **Minho Kim (minos@ecnebiz.com), ECNE Inc.
 - ***Namgyu Kim (ngkim@kookmin.ac.kr), Graduate School of Business IT, Kookmin University
 - Received: 2026. 04. 21, Revised: 2026. 05. 22, Accepted: 2026. 06. 03.

I. Introduction

딥러닝 기술의 발전과 함께 온라인 콘텐츠 시장이 급속히 성장하면서, 사용자가 방대한 양의 콘텐츠 중에서 자신의 취향에 맞는 정보를 효율적으로 탐색하는 것이 점점 더 어려워지고 있다. 이에 따라 사용자의 과거 행동 패턴을 분석하여 개인화된 콘텐츠를 제공하는 추천 시스템의 중요성이 높아지고 있으며[1], 특히 전자상거래, 스트리밍 서비스, 온라인 플랫폼 등 다양한 분야에서 추천 시스템은 사용자 경험 향상과 플랫폼 경쟁력 확보의 핵심 기술로 자리매김하고 있다. 그러나 신규 IP(Intellectual Property) 콘텐츠 환경에서는 출시 초기에 사용자-아이템 상호작용 데이터가 절대적으로 부족한 Cold-start 문제가 발생하여 [2, 3], 충분한 학습 데이터를 확보하지 못해 정확한 추천을 제공하기 어렵다는 근본적인 한계가 존재한다.

이러한 환경에서 순차 추천(Sequential Recommendation) 방식은 사용자의 과거 소비 이력을 시퀀스로 모델링하여 다음 아이템을 예측하는 접근법으로 주목받고 있다[4]. 대표적으로 SASRec[5]은 Transformer의 Self-Attention 메커니즘을 활용하여 사용자의 장단기 선호를 효과적으로 모델링하였으나, 아이템 ID 임베딩에 의존하는 구조상 신규 아이템에 대한 추천이 어렵다는 한계를 지닌다.

이를 극복하기 위해 아이템의 텍스트 속성 정보를 직접 활용하는 텍스트 기반 추천 모델이 제안되었다. 구체적으로 RecFormer[6]는 긴 문서 처리에 특화된 Longformer[7] 구조를 기반으로, 아이템의 제목·카테고리 등 텍스트 속성을 키-밸류(key-value) 쌍 형태의 문장으로 구성하고 사용자의 전체 상호작용 시퀀스를 하나의 긴 문서로 처리한다. 이를 통해 신규 아이템에 대한 추천이 가능하며, 저자원 및 Cold-start 환경에서도 높은 성능을 보이는 것으로 알려져 있다. 그러나 RecFormer 역시 학습 데이터가 희소한 환경에서는 충분한 패턴 학습이 어렵다는 한계가 있다[8]. 즉, 신규 IP 추천과 같이 사용자 수가 극히 적거나 상호작용 이력이 제한적인 상황에서는 절대적인 학습 샘플 수가 부족하여, 모델이 다양한 사용자 선호 패턴을 학습하지 못해 추천 성능이 저하될 수 있다.

본 논문에서는 이러한 데이터 희소 문제를 완화하기 위해, RecFormer에 Sliding Window[9] 기반 학습 데이터 증강(Data Augmentation) 기법을 적용하는 방법을 제안한다. 제안 방법은 각 사용자의 상호작용 시퀀스를 고정 크기의 윈도우로 분할하여 다수의 부분 시퀀스를 생성함으로써, 하나의 사용자 시퀀스로부터 다수의 학습 샘플을 확보하는 것을 목표로 한다. 이를 통해 절대적인 학습 샘플

수를 효과적으로 증가시킬 뿐만 아니라, 시퀀스 내 다양한 위치의 국소 패턴을 학습함으로써 모델의 일반화 능력을 향상시키고자 한다. 특히 본 연구는 모델 구조의 변경 없이 학습 데이터 구성 방식만을 개선하는 접근법을 취하고 있어, 기존 RecFormer 구조를 그대로 유지하면서도 데이터 희소 환경에서의 추천 성능을 향상시킬 수 있을 것으로 기대한다.

본 논문의 이후 구성은 다음과 같다. 2장에서는 순차 추천, 텍스트 기반 추천, 데이터 증강에 관한 선행 연구를 요약하고, 3장에서는 본 연구에서 제안하는 슬라이딩 윈도우 기반 데이터 증강을 통한 RecFormer 성능 향상 방안 방법론을 상세히 설명한다. 4장에서는 실험 설정 및 성능 평가 결과를 분석한다. 마지막으로 5장에서는 본 연구의 기여와 한계 및 향후 연구 방향을 제시한다.

II. Preliminaries

1. Sequence-based Recommendation

순차 추천은 사용자의 과거 상호작용 이력을 시간 순서로 모델링하여 다음에 소비할 아이템을 예측하는 방식이다. 초기에는 행렬 분해(Matrix Factorization)와 Markov Chain을 결합하여 사용자의 개인화된 순차적 행동 패턴을 모델링하는 방법론이 활용되었다[10]. 이러한 접근법은 사용자의 단기 순차 패턴을 포착하는 데 효과적이었으나, 고차원의 순차적 의존 관계를 모델링하는 데 한계가 있었다. 이후 딥러닝의 발전과 함께 RNN 기반 모델인 GRU4Rec[11]이 제안되어 세션 기반 추천에서 순차 패턴 모델링 능력을 향상시켰다.

Transformer 아키텍처의 등장으로 순차 추천의 패러다임은 크게 변화하였다. SASRec[5]은 단방향 Self-Attention 메커니즘을 활용하여 사용자 시퀀스 내 모든 아이템 간의 의존 관계를 포착함으로써, 장기 및 단기 선호를 효과적으로 모델링하였다. BERT4Rec[12]은 BERT의 양방향 Transformer 구조를 순차 추천에 적용하고, Cloze task 방식의 마스킹 학습을 통해 시퀀스 내 양방향 문맥 정보를 활용하여 성능을 향상시켰다.

그러나 이러한 모델들은 공통적으로 아이템 ID 임베딩에 의존하기 때문에, 학습 과정에서 등장하지 않은 신규 아이템에 대해서는 추천이 불가능하다는 근본적인 한계를 지닌다. 또한 아이템 ID 기반 임베딩은 도메인 간 전이 학습이 어렵고, 학습 데이터가 희소한 환경에서는 충분한 임베딩 학습이 이루어지지 않아 추천 성능이 크게 저하될 수 있다는 한계를 갖는다.

2. Text-based Recommendation

텍스트 기반 추천은 아이템 ID 대신 아이템의 텍스트 속성 정보를 활용하여 사용자 선호를 모델링하는 방식으로, ID 기반 모델의 Cold-start 한계를 극복하고자 하는 연구 흐름에서 발전하였다.

ZESRec[13]은 아이템의 자연어 설명을 범용 식별자로 활용하여, 학습 도메인과 테스트 도메인 간 사용자와 아이템이 전혀 겹치지 않는 Zero-shot 추천 환경에서도 일반화 가능한 사용자 행동 패턴을 학습하였다. UniSRec[14]은 아이템의 텍스트 설명을 사전학습 언어 모델로 인코딩하고, MoE(Mixture-of-Experts) 기반 어댑터를 통해 다양한 도메인 간 전이 학습이 가능한 범용 시퀀스 표현을 학습하였다. 이를 통해 아이템 ID 없이도 다양한 추천 시나리오에 적용 가능한 보편적 표현 학습을 달성하였다.

RecFormer[6]는 Longformer[7]를 기반으로 아이템의 텍스트 속성(제목, 카테고리 등)을 키패럴 쌍 형태의 문장으로 구성하고, 사용자의 전체 상호작용 시퀀스를 하나의 긴 문서로 처리하여 추천을 수행한다. 사전학습과 2단계 파인튜닝(Fine-Tuning)을 통해 언어 이해와 추천 학습을 결합하며, 특히 저자원 및 Cold-start 환경에서 뛰어난 성능을 보인다. 그러나 RecFormer 역시 학습 데이터 자체가 절대적으로 부족한 환경에서는 다양한 패턴을 충분히 학습하기 어렵다는 한계가 존재한다.

3. Data Augmentation in Recommendation

데이터 증강은 부족한 학습 데이터를 인위적으로 확장하여 모델의 일반화 능력을 향상시키는 기법으로, 추천 시스템 분야에서도 활발히 연구되고 있다[15].

CL4SRec[16]은 순차 추천에 대조 학습(Contrastive Learning)을 적용한 연구로, 아이템 크로핑(Cropping), 마스킹(Masking), 재정렬(Reordering)의 세 가지 랜덤 증강 기법을 통해 동일 시퀀스의 서로 다른 뷰를 생성한다. 생성된 두 가지 뷰 간의 표현 일치도를 최대화하는 대조 학습 목적 함수를 통해 더 풍부한 사용자 표현을 학습함으로써, 데이터 희소 환경에서의 추천 성능을 향상시켰다. CoSeRec[17]은 아이템 간 상관관계를 고려한 대체(Substitute) 및 삽입(Insert) 연산자를 도입하여, 의미적으로 일관성 있는 증강 샘플을 생성하는 방식을 제안하였다. 이를 통해 더 정보성 높은 학습 뷰를 생성하고 추천 성능을 향상시켰다.

한편 Caser[9]는 합성곱 신경망(Convolutional Neural Network, CNN) 기반 순차 추천 모델로, 각 사용자의 시퀀스에 Sliding Window를 적용하여 다수의 학습 샘플을

생성하는 방식을 제안하였다. 구체적으로 크기 $L+T$ 의 윈도우를 사용자 시퀀스 위에 이동시키며, 각 윈도우에서 이전 L 개 아이템을 입력으로, 다음 T 개 아이템을 타겟으로 하는 학습 인스턴스를 생성한다. 본 논문에서는 이러한 Sliding Window 방식을 텍스트 기반 순차 추천 모델인 RecFormer의 파인튜닝 단계에 적용함으로써, 기존의 단일 샘플 생성 방식의 한계를 극복하고 데이터 희소 환경에서의 추천 성능을 향상시키고자 한다. 이는 CL4SRec[16]의 Cropping 기법이 랜덤한 위치에서 시퀀스 일부를 무작위로 선택하여 대조 학습을 위한 뷰를 생성하는 방식과 달리, 고정된 Window Size와 Stride에 따라 사용자 시퀀스를 체계적으로 분할하여 파인튜닝에 활용 가능한 학습 샘플을 확장한다는 점에서 차별성을 지닌다.

III. The Proposed Scheme

1. Research Process

본 장에서는 데이터 희소 환경에서 RecFormer의 추천 성능을 향상시키기 위해 Sliding Window 기반 데이터 증강을 적용하는 전체 방법론을 각 단계별로 소개한다. 제안 방법론은 총 두 단계로 구성되며, 전체적인 과정은 Fig. 1과 같다.

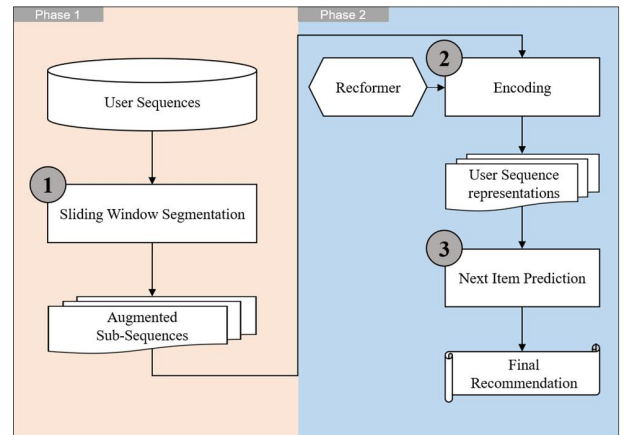


Fig. 1. Overall Research Process

Phase 1에서는 사용자 시퀀스를 기반으로 학습 데이터를 확장하기 위한 Sliding Window 기반 시퀀스 분할을 수행한다. 기존의 시퀀스 추천 모델은 하나의 사용자 상호작용 시퀀스로부터 단일 학습 샘플을 생성하는 반면, 본 연구에서는 일정한 Window Size와 Stride를 기반으로 시퀀스를 분할하여 다수의 부분 시퀀스를 생성한다(①). 이를 통해 데이터가 부족한 환경에서도 다양한 사용자 행동

패턴을 학습할 수 있도록 한다. 생성된 부분 시퀀스는 Augmented Sub-Sequences로 구성되며, 이후 모델 학습 단계의 입력으로 사용된다.

Phase 2에서는 증강된 시퀀스를 활용하여 Recformer 기반 추천 모델을 학습하고, 최종 추천 결과를 생성한다. 먼저, 각 부분 시퀀스는 Recformer 모델에 입력되어 Encoding 과정을 거치며, 이를 통해 시퀀스의 문맥 정보를 반영한 사용자 시퀀스 표현(User Sequence Representations)이 생성된다(②). 이후 생성된 시퀀스 표현은 Next Item Prediction 단계에 입력되어, 사용자의 다음 상호작용 아이템을 예측한다(③). 해당 과정에서는 시퀀스 표현을 기반으로 후보 아이템에 대한 점수를 산출하고, 이를 바탕으로 최종 추천 결과를 도출한다. 최종적으로 상위 K개의 아이템을 선택하여 추천한다.

각 단계에 대한 세부 프로세스는 이후 절에서 상세히 설명하며, 실제 데이터를 적용한 제안 방법론의 성능 평가 결과는 4장에서 소개한다.

2. RecFormer Architecture

본 절에서는 Fig. 1의 Phase 2에서, 사용자의 상호작용 시퀀스를 인코딩하여 사용자 시퀀스 표현을 생성하고(②), 이를 기반으로 다음 아이템을 예측하는 RecFormer의 구조를 소개한다(③).

RecFormer[6]는 긴 문서 처리에 특화된 Longformer[7]를 기반으로 설계된 텍스트 기반 순차 추천 모델로, 기존 ID 기반 순차 추천 모델과 달리 아이템의 텍스트 속성 정보를 직접 활용하여 사용자 선호도를 모델링함으로써 신규 아이템에 대한 추천이 가능하다. RecFormer는 각 아이템의 제목(Title), 카테고리(Category), 브랜드(Brand) 등의 텍스트 속성을 키-밸류 쌍 형태로 평탄화(Flattening)한 아이템 문장으로 표현한다. 사용자의 상호작용 시퀀스 $s = \{i_1, i_2, \dots, i_n\}$ 에 대해 각 아이템의 텍스트 표현을 역순으로 연결하고, 맨 앞에 [CLS] 토큰을 추가하여 하나의 긴 입력 시퀀스 X 를 구성한다. 이때 역순으로 배열하는 이유는 최근 상호작용 아이템이 다음 아이템 예측에 더 중요하게 사용되기 때문이다.

RecFormer의 임베딩 레이어는 토큰 임베딩(Token Embedding), 토큰 위치 임베딩(Token Position Embedding), 토큰 타입 임베딩(Token Type Embedding), 아이템 위치 임베딩(Item Position Embedding)의 네 가지 임베딩을 결합하여 입력 표현을 구성한다. 이를 통해 언어 모델의 텍스트 이해 능력과 순차 추천 모델의 시퀀스 패턴 학습 능력을 동시에 활용할 수 있다.

인코딩 과정에서는 Longformer의 지역 윈도우 어텐션(Local Windowed Attention)과 전역 어텐션(Global Attention)을 조합하여 긴 입력 시퀀스를 효율적으로 처리한다. [CLS] 토큰에는 전역 어텐션이 적용되어 시퀀스 전체와 상호작용하며, 해당 토큰의 최종 은닉 상태 $h[\text{CLS}]$ 가 사용자 시퀀스 표현으로 사용된다. 이렇게 생성된 사용자 시퀀스 표현 $h[\text{CLS}]$ 와 후보 아이템 표현 h_i 간의 코사인 유사도를 계산하여, 유사도가 가장 높은 아이템을 다음 아이템으로 예측한다(③).

RecFormer는 사전학습(Pre-Training)과 2단계 파인튜닝(Two-Stage Fine-Tuning)을 통해 학습된다. 사전학습 단계에서는 마스킹 언어 모델링과 아이템 간 대조 학습을 결합하여 언어 이해와 추천 패턴을 동시에 학습하며, 이후 파인튜닝 단계에서 목표 도메인의 상호작용 데이터를 활용하여 특정 추천 태스크에 적응시킨다. 본 연구에서는 3.3절에서 소개하는 Sliding Window 기반 증강 데이터를 활용하여 파인튜닝을 수행한다.

3. Sliding Window-based Data Augmentation

본 절에서는 Fig. 1의 Phase 1에서, 사용자의 상호작용 시퀀스에 Sliding Window를 적용하여 증강된 학습 샘플을 생성하는 과정을 설명한다(①).

기존 RecFormer의 학습 방식에서는 사용자 시퀀스 $s = \{i_1, i_2, \dots, i_n\}$ 에 대해 입력 $\{i_1, \dots, i_{n-1}\}$ 과 타겟 i_n 의 단일 학습 샘플만이 생성된다. 이러한 방식은 데이터가 충분한 환경에서는 큰 문제가 되지 않지만, 사용자 수가 적거나 상호작용 이력이 제한적인 데이터 희소 환경에서는 절대적인 학습 샘플 수가 부족하여 모델이 다양한 패턴을 학습하기 어렵다는 한계를 가진다.

본 연구에서는 이러한 한계를 극복하기 위해 Sliding Window 기반 시퀀스 분할 기법을 적용한다. Sliding Window는 고정 크기의 윈도우를 사용자 시퀀스 위에 일정 간격으로 이동시키며 다수의 부분 시퀀스를 생성하는 방법으로, 각 학습 샘플은 길이 w 의 입력 시퀀스와 그에 대응하는 다음 아이템을 타겟으로 하는 학습 쌍으로 구성된다(Fig. 2).

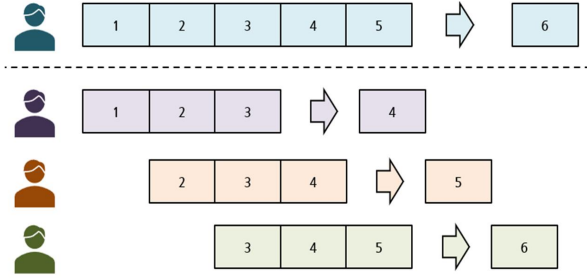


Fig. 2. Example of Sliding Window Segmentation

예를 들어, 하나의 사용자 시퀀스 [1, 2, 3, 4, 5, 6]에 대해 기존 방식에서는 [1, 2, 3, 4, 5] → 6의 단일 학습 샘플만 생성된다. 하지만 Fig. 2와 같이 윈도우 크기를 3, 이동 간격을 1로 설정하여 Sliding Window를 적용할 경우 [1, 2, 3] → 4, [2, 3, 4] → 5, [3, 4, 5] → 6과 같은 세 개의 학습 샘플이 생성된다. 이를 통해 하나의 시퀀스로부터 다수의 학습 데이터를 확보할 수 있다.

제안하는 Sliding Window 기반 데이터 증강은 RecFormer의 파인튜닝 단계에 직접 적용된다. 생성된 부분 시퀀스들은 각각 독립적인 학습 샘플로 모델에 입력되며, 이를 통해 RecFormer는 다양한 위치의 국소 시퀀스 패턴을 학습할 수 있다. 특히 기존 방식에서는 학습되지 않던 시퀀스 중간 구간의 패턴까지 학습에 활용할 수 있다는 점에서, 모델의 표현 학습 범위를 확장하는 효과를 가진다.

IV. Experiment

1. Experiment Overview

본 장에서는 제안 방법론의 성능을 평가하기 위해 수행한 실험 과정 및 결과를 소개한다.

실험에는 Amazon Scientific Review 데이터셋을 활용하였다. 해당 데이터셋은 과학 관련 상품에 대한 사용자 리뷰 및 평점 데이터로 구성되며, 전처리 후 총 11,039개의 사용자 시퀀스와 5,327개의 고유 아이템을 포함한다. 각 아이템은 제목, 카테고리, 브랜드 등의 텍스트 속성 정보를 포함하며, RecFormer의 입력 형식에 맞게 키-밸류 쌍으로 구성되었다. 데이터 희소 환경을 시뮬레이션하기 위해 전체 학습 데이터의 1%와 0.5%만을 사용하는 제한적 환경에서 실험을 수행하였으며, 검증 및 테스트 데이터는 Leave-One-Out 방식을 적용하였다. 학습은 총 4 epoch 동안 진행하였으며, 배치 크기는 4, 학습률은 $5e-5$ 로 설정하였다. 또한 Window Size는 3, Stride는 1로 설정하였다. Window Size 3은 최소한의 순차 문맥을 유지

하면서도 충분한 수의 부분 시퀀스를 생성하기 위한 설정이며, Stride 1은 가능한 연속 부분 시퀀스를 학습 샘플로 활용하여 데이터 증강 효과를 극대화하기 위한 설정이다. 본 실험은 Python 3.12 환경에서 수행되었으며, 구체적인 실험 환경은 Table 1과 같다.

Table 1. System Environment

HW	CPU	24 Core
	GPU	NVIDIA A100 40GB X2
	Memory	192 GiB
SW	OS	Ubuntu 24.04.1 LTS
	Python	3.12
	Pytorch	2.6

비교 모델로는 Sliding Window 증강을 적용하지 않은 기존 RecFormer를 사용하였으며, 제안 방법과의 공정한 비교를 위해 모델 구조 및 학습 설정은 동일하게 유지하였다. 전체적인 성능 평가 프로세스는 Fig. 3과 같다.

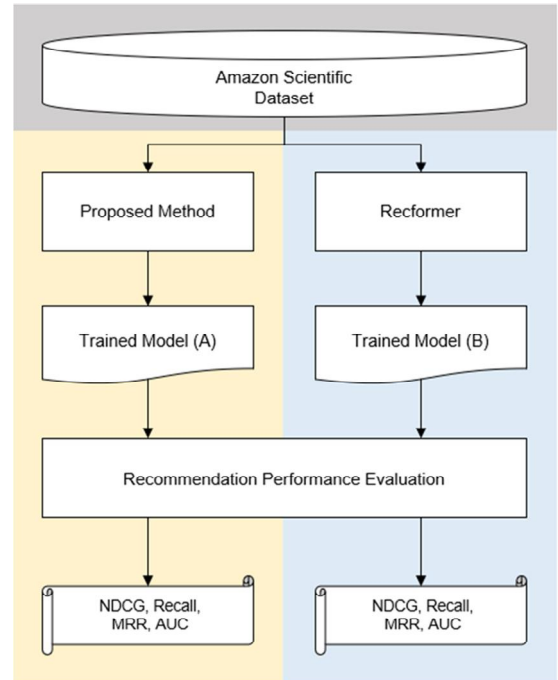


Fig. 3. Overall Process of Performance Evaluation

Fig. 3과 같이, 동일한 Amazon Scientific 데이터셋을 기준으로 제안 방법론(A)과 기존 RecFormer(B)를 각각 학습한 후, 추천 성능 평가를 수행하여 NDCG, Recall, MRR, AUC 등 다양한 평가 지표를 활용하여 두 모델의 성능을 종합적으로 비교하였다.

2. Evaluation Metrics

본 절에서는 제안 방법론의 추천 성능을 평가하기 위해 사용한 평가 지표를 소개한다.

본 실험에서는 순차 추천 시스템의 성능 평가에 널리 활용되는 Recall@k, NDCG@k, MRR, AUC를 평가 지표로 사용하였으며, k는 10과 50으로 설정하였다.

Recall@k는 모델이 실제 정답 아이템을 얼마나 놓치지 않고 추천 목록 안에 포함시키는지 평가하는 지표이다 (식 1). 따라서 Recall@k 값이 높을수록 모델이 사용자의 실제 선호 아이템을 더 잘 포착하고 있음을 의미한다. 특히 추천 목록의 상위 k개 안에 정답 아이템이 포함되는 비율을 측정하므로, 모델의 전반적인 추천 포착 능력을 평가하는 데 적합하다.

$$Recall@k = \frac{1}{|U|} \sum_{u \in U} \frac{|R_u^{(k)} \cap G_u|}{|G_u|} \dots (1)$$

NDCG@k(Normalized Discounted Cumulative Gain)는 추천된 아이템의 순위를 고려하는 지표로, 상위에 위치한 정답 아이템에 더 높은 가중치를 부여한다. 먼저 DCG(Discounted Cumulative Gain)를 계산한 후, 이를 이상적인 순서의 IDCG(Ideal DCG)로 정규화하여 다음과 같이 산출한다(식 2).

$$NDCG@k = \frac{DCG@k}{IDCG@k}, DCG@k = \sum_{i=1}^k \frac{rel_i}{\log_2(i+1)} \dots (2)$$

여기서 rel_i 는 순위 i에 위치한 아이템의 관련도를 의미한다. NDCG@k는 단순히 정답 아이템의 포함 여부만이 아니라, 해당 아이템이 추천 목록의 어느 위치에 배치되었는지를 함께 반영한다. 따라서 NDCG@k 값이 높을수록 정답 아이템이 추천 목록의 상위에 배치되어 있음을 의미하며, 결과적으로 추천의 순위 품질이 우수하다고 해석할 수 있다. 즉, 동일하게 정답 아이템을 포함하더라도 더 앞 순위에 제시한 모델이 더 높은 NDCG@k 값을 나타낸다.

MRR(Mean Reciprocal Rank)은 각 사용자에게 대해 정답 아이템이 추천 목록에서 처음 나타나는 순위의 역수를 평균한 값으로 정의되며, 다음과 같이 계산된다(식 3).

$$MRR@k = \frac{1}{|U|} \sum_{u=1}^U \frac{1}{rank_u} \dots (3)$$

여기서 $rank_u$ 는 사용자 u의 정답 아이템이 추천 목록

에서 나타나는 순위를 의미하므로, MRR은 정답 아이템 얼마나 빠르게 추천되는지를 평가하는 지표로 볼 수 있다. 따라서 MRR 값이 높을수록 정답 아이템이 추천 목록의 더 상위에 위치함을 의미하며, 사용자가 추천 결과를 확인할 때 더 적은 탐색 비용으로 원하는 아이템에 도달할 수 있음을 나타낸다. MRR은 하나의 핵심 정답 아이템을 얼마나 높은 순위에 노출시키는지 중요한 환경에서 유용한 지표이다.

AUC(Area Under the Curve)는 모델이 정답 아이템을 비정답 아이템보다 더 높은 순위에 배치하는 능력을 평가하는 지표이다. 이는 추천 결과 전체에서 모델의 판별력을 측정하는 값으로, 일반적으로 0.5는 무작위 수준의 성능을 의미하고 1에 가까울수록 정답 아이템과 비정답 아이템을 더 잘 구분함을 의미한다. 따라서 AUC 값이 높을수록 모델이 전반적으로 더 안정적이고 일관되게 올바른 순위를 형성하고 있다고 해석할 수 있다.

이러한 다양한 평가 지표를 함께 활용함으로써, 본 연구에서는 추천 정확도, 순위 품질, 상위 노출 성능, 그리고 전반적인 판별 능력 등 여러 측면에서 제안 방법론의 성능을 종합적으로 분석하였다.

3. Performance Evaluation

본 절에서는 제안 방법론과 기존 RecFormer의 추천 성능 비교 결과를 분석한다. 학습 데이터 비율 1% 및 0.5% 환경에서의 성능 비교 결과는 각각 Table 2 및 Table 3에 요약되어 있다.

Table 2와 같이, 제안 방법론은 1% 학습 데이터 환경에서 NDCG@10, Recall@10, NDCG@50, Recall@50, MRR, AUC 등 모든 지표에서 기존 RecFormer 대비 우수한 성능을 보였다. 특히 NDCG@10은 0.0939를 기록하여 기존 RecFormer의 0.0922 대비 약 1.8% 향상되었으며, AUC 또한 0.6991로 기존 모델의 0.6924보다 높은 값을 나타냈다. 이는 Sliding Window 기반 데이터 증강을 통해 학습 샘플 수가 증가함에 따라, 모델이 보다 다양한 사용자 행동 패턴을 학습할 수 있었기 때문으로 해석된다.

Table 3과 같이, 0.5% 학습 데이터 환경에서는, 데이터 희소성이 더욱 심화된 상황에서도 제안 방법론이 안정적인 성능 향상을 보였다. NDCG@10은 0.0927로 기존 RecFormer의 0.0889 대비 약 4.3% 향상되었으며, NDCG@50, MRR, AUC 지표에서도 일관된 개선이 확인되었다. 반면 Recall@10과 Recall@50은 기존 RecFormer 대비 소폭 감소하였다. 이는 제안 방법론이 더 많은 정답 아이템을 상위 k개 목록에 포함시키는 포착 범위 측면에

Table 2. Performance Comparison (Train Size 1%)

Method	NDCG@10	Recall@10	NDCG@50	Recall@50	MRR	AUC
Proposed Method	<u>0.0939</u>	<u>0.1306</u>	<u>0.1063</u>	<u>0.1869</u>	<u>0.0865</u>	<u>0.6991</u>
RecFormer	0.0922	0.1291	0.1048	0.1866	0.0847	0.6924

Table 3. Performance Comparison (Train Size 0.5%)

Method	NDCG@10	Recall@10	NDCG@50	Recall@50	MRR	AUC
Proposed Method	<u>0.0927</u>	0.1257	<u>0.1048</u>	0.1813	<u>0.0865</u>	<u>0.6897</u>
RecFormer	0.0889	<u>0.1282</u>	0.1006	<u>0.1814</u>	0.0805	0.6875

서는 일부 제한을 보였으나, NDCG@10과 MRR이 향상된 결과를 고려할 때 정답 아이템이 추천 목록 내에서 더 높은 순위에 배치되는 방향으로 순위 정렬 품질이 개선되었음을 의미한다. 즉, 제안 방법론은 데이터 희소 환경에서 추천 후보의 포괄성보다는 상위 순위 정렬 품질을 향상시키는 경향을 보였으며, 이는 Recall과 NDCG 간의 부분적인 trade-off로 해석할 수 있다.

데이터 증강 기법이 데이터 희소 환경에서의 순차 추천 성능 향상에 효과적임을 확인할 수 있다.

V. Conclusions

디지털 콘텐츠 시장의 급격한 성장으로 추천 시스템의 중요성이 강조되고 있으나, 신규 서비스 초기와 같이 학습 데이터가 충분하지 않은 환경에서는 추천 성능이 크게 저하되는 문제가 발생한다. 이에 본 연구에서는 텍스트 기반 순차 추천 모델인 RecFormer의 파인튜닝 과정에 Sliding Window 기반 데이터 증강을 적용하여 데이터 희소 환경에서의 추천 성능을 향상시키는 방법론을 제안하였다. Amazon Scientific Review 데이터셋을 활용한 실험 결과, 제안 방법론은 1% 및 0.5% 학습 데이터 환경에서 기존 RecFormer 대비 대부분의 평가 지표에서 향상된 성능을 달성하였으며, 특히 학습 데이터가 더욱 제한적인 환경일수록 성능 향상 효과가 더욱 두드러지게 나타났다.

본 연구는 텍스트 기반 순차 추천 모델에 데이터 증강 기법을 접목한 새로운 학습 프레임워크를 제안하였다는 점에서 학술적으로 의미가 있다. 기존 연구들이 모델 구조 개선에 초점을 맞추었던 것과 달리, 본 연구는 학습 데이터 구성 방식의 개선을 통해 데이터 희소 환경에서의 추천 성능을 향상시킬 수 있음을 실험적으로 확인하였다. 실용적 측면에서도 제안 방법론은 모델 구조의 변경 없이 학습 데이터 구성 방식만을 개선하는 접근법을 취하고 있어 구현이 간단하고 기존 시스템에 쉽게 적용할 수 있다는 장점을 지닌다. 특히 신규 콘텐츠 플랫폼의 초기 서비스 환경, 데이터 수집 비용이 높은 전문 분야 추천 시스템, 또는 사용자 수가 제한적인 소규모 서비스 환경 등 다양한 데이터 희소 상황에서 실질적으로 활용 가능한 방법론으로서의 의미를 갖는다.

본 연구에는 몇 가지 한계가 존재하며, 이를 바탕으로

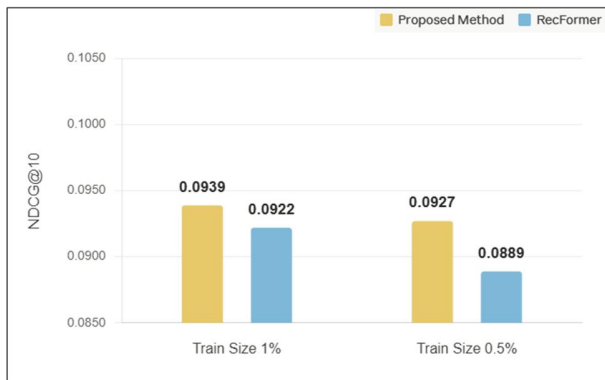


Fig. 4. NDCG@10 Performance Comparison by Train Size

특히 주목할 점은 데이터 희소성이 증가할수록 성능 향상 폭이 더욱 확대된다는 점이다. Fig. 4에서와 같이 1% 환경 대비 0.5% 환경에서 NDCG@10 기준 성능 개선 폭이 더 크게 나타났으며, 이는 Sliding Window 기반 데이터 증강이 제한된 데이터 환경에서 더욱 효과적으로 작동함을 의미한다. 즉, 단일 시퀀스로부터 다수의 학습 샘플을 생성함으로써 모델이 다양한 시퀀스 패턴을 학습할 수 있게 되어, 데이터 부족으로 인한 성능 저하를 완화하는 효과를 가진다.

종합적으로, 제안 방법론은 데이터 희소 환경에서 학습 샘플 수를 효과적으로 확장함으로써 기존 RecFormer 대비 전반적인 추천 성능을 향상시켰다. 특히 극도로 제한된 데이터 환경에서 성능 개선 효과가 더욱 두드러지게 나타났다. 본 연구에서 제안한 Sliding Window 기반

향후 연구 방향을 제시하고자 한다. 본 연구에서는 window size와 stride를 고정하여 실험을 수행하였으나, 향후 연구에서는 다양한 하이퍼파라미터 조합에 따른 성능 변화를 체계적으로 분석하여 최적의 설정을 탐색할 필요가 있다. 또한 단일 도메인 데이터셋만을 활용하였으므로, 향후 다양한 도메인의 데이터셋에 대한 추가 실험을 통해 제안 방법론의 범용성을 검증할 필요가 있다. 추가적으로 본 연구에서는 제안 방법론의 추천 성능 개선 효과를 중심으로 분석하였으며, 학습 및 평가 과정에 소요되는 시간과 GPU 메모리 사용량에 대한 정량적 비교는 수행하지 못하였다. 제안 방법론은 Sliding Window 적용으로 학습 샘플 수가 증가하므로, 성능 향상과 함께 계산 비용이 증가할 가능성이 있다. 향후 연구에서는 학습 시간, 평가 시간, 메모리 사용량 등을 포함한 계산 효율성 분석을 추가적으로 수행할 필요가 있다. 마지막으로, 다른 데이터 증강 기법과의 결합을 통해 추가적인 성능 향상 가능성을 탐색하는 것도 의미 있는 연구 방향이 될 것이다.

ACKNOWLEDGEMENT

This work was supported by the Starting growth Technological R&D Program (TIPS Program, (No. RS-2024-00511799)) funded by the Ministry of SMEs and Startups(MSS, Korea) in 2024

REFERENCES

- [1] S. Zhang, L. Yao, A. Sun, and Y. Tay, "Deep Learning based Recommender System: A Survey and New Perspectives," *ACM Computing Surveys*, Vol. 52, No. 1, pp. 1-38, September 2019. DOI: 10.1145/3285029
- [2] H. Yuan and A. A. Hernandez, "User Cold Start Problem in Recommendation Systems: A Systematic Review," *IEEE Access*, Vol. 11, pp. 136958-136977, December 2023. DOI: 10.1109/ACCESS.2023.3338705
- [3] Y. Liu, "The Solution of Cold Start Problem in Recommender System Introduction," *Theory and Practice of Science and Technology*, Vol. 4, No. 1, pp. 48-54, 2023.
- [4] P. Wei, H. Shu, J. Gan, X. Deng, Y. Liu, W. Sun, T. Chen, C. Hu, Z. Hu, Y. Deng, W. Qin, and Z. Li, "Sequential Recommendation System Based on Deep Learning: A Survey," *Electronics*, Vol. 14, No. 11, p. 2134, May 2025. DOI: 10.3390/electronics14112134
- [5] W.-C. Kang and J. McAuley, "Self-Attentive Sequential Recommendation," 2018 IEEE International Conference on Data Mining (ICDM), pp. 197-206, Singapore, Singapore, November 2018. DOI: 10.1109/ICDM.2018.00035
- [6] J. Li, M. Wang, J. Li, J. Fu, X. Shen, J. Shang, and J. McAuley, "Text Is All You Need: Learning Language Representations for Sequential Recommendation," *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD '23)*, pp. 1258-1267, Long Beach, CA, USA, August 2023. DOI: 10.1145/3580305.3599519
- [7] I. Beltagy, M. E. Peters, and A. Cohan, "Longformer: The Long-Document Transformer," *arXiv preprint arXiv:2004.05150*, April 2020. DOI: 10.48550/arXiv.2004.05150
- [8] J. Li, T. Zhao, J. Li, J. Chan, C. Faloutsos, G. Karypis, S.-M. Pantel, and J. McAuley, "Coarse-to-Fine Sparse Sequential Recommendation," *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '22)*, pp. 2082-2086, Madrid, Spain, July 2022. DOI: 10.1145/3477495.3531732
- [9] J. Tang and K. Wang, "Personalized Top-N Sequential Recommendation via Convolutional Sequence Embedding," *Proceedings of the Eleventh ACM International Conference on Web Search and Data Mining (WSDM '18)*, pp. 565-573, Marina Del Rey, CA, USA, February 2018. DOI: 10.1145/3159652.3159656
- [10] S. Rendle, C. Freudenthaler, and L. Schmidt-Thieme, "Factorizing Personalized Markov Chains for Next-Basket Recommendation," *Proceedings of the 19th International Conference on World Wide Web (WWW '10)*, pp. 811-820, Raleigh, NC, USA, April 2010. DOI: 10.1145/1772690.1772773
- [11] B. Hidasi, A. Karatzoglou, L. Baltrunas, and D. Tikk, "Session-based Recommendations with Recurrent Neural Networks," *arXiv preprint arXiv:1511.06939*, November 2015. DOI: 10.48550/arXiv.1511.06939
- [12] F. Sun, J. Liu, J. Wu, C. Pei, X. Lin, W. Ou, and P. Jiang, "BERT4Rec: Sequential Recommendation with Bidirectional Encoder Representations from Transformer," *Proceedings of the 28th ACM International Conference on Information and Knowledge Management (CIKM '19)*, pp. 1441-1450, Beijing, China, November 2019. DOI: 10.1145/3357384.3357895
- [13] H. Ding, Y. Ma, A. Deoras, Y. Wang, and H. Wang, "Zero-Shot Recommender Systems," *arXiv preprint arXiv:2105.08318*, May 2021. DOI: 10.48550/arXiv.2105.08318
- [14] Y. Hou, S. Mu, W. X. Zhao, Y. Li, B. Ding, and J. Wen, "Towards Universal Sequence Representation Learning for Recommender Systems," *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD '22)*, pp. 585-593, Washington, DC, USA, August 2022. DOI: 10.1145/3534678.3539381

- [15] Y. Dang, E. Yang, Y. Liu, G. Guo, L. Jiang, J. Zhao, and X. Wang, "Data Augmentation for Sequential Recommendation: A Survey," arXiv preprint arXiv:2409.13545, September 2024. DOI: 10.48550/arXiv.2409.13545
- [16] X. Xie, F. Sun, Z. Liu, S. Wu, J. Gao, J. Zhang, B. Ding, and B. Cui, "Contrastive Learning for Sequential Recommendation," 2022 IEEE 38th International Conference on Data Engineering (ICDE), pp. 1259-1273, Kuala Lumpur, Malaysia, May 2022. DOI: 10.1109/ICDE53745.2022.00099
- [17] Z. Liu, Y. Chen, J. Li, P. S. Yu, J. McAuley, and C. Xiong, "Contrastive Self-supervised Sequential Recommendation with Robust Augmentation," arXiv preprint arXiv:2108.06479, August 2021. DOI: 10.48550/arXiv.2108.06479

Authors



Sungho Choi received the B.S. degree in Computer Science from Sangmyung University and is currently enrolled in the Graduate School of Business IT, Kookmin University. He is interested in deep learning,

large language models, and sentence embeddings.



Minho Kim received the B.S. degree in Management Information Systems from the School of Business IT, Kookmin University. Minh Kim is the President of ECNE, where he manages a specialized content platform.

He has extensive experience in conducting multiple LLM-related research projects for the Ministry of SMEs and Startups (MSS). His current focus is on the practical application of LLMs and digital content ecosystems.



Namgyu Kim received the B.S. in Computer Engineering from Seoul National University in 1998, M.S. and Ph.D. degrees in Management Engineering from KAIST, Korea, in 2000 and 2007, respectively.

Dr. Kim joined the faculty of the School of Management Information Systems at Kookmin University, Seoul, Korea, in 2007. He served as the Dean of the Graduate School of Business IT at Kookmin University and is currently a professor at the Business IT. He is interested in LLM, text mining, and deep learning.