

저작권 기술 트렌드



COPYRIGHT TECH TREND



한국저작권위원회
KOREA COPYRIGHT COMMISSION

AI 기반 게임 이용자 행동 데이터 분석 기술

뉴스 브리프

게임 이용자가 게임과 상호작용하며 남기는 행동 기록에는 이용자의 참여도와 선호가 담겨 있어 이를 분석해 이용자의 향후 행동을 예측하려는 연구가 이어지고 있다. 이러한 행동 데이터 분석에는 다양한 접근이 있으며 본 보고서는 그중 서로 다른 데이터 상황에 대응하는 두 건의 연구를 살펴본다. 첫 번째 연구는 한 게임 안에 행동 이력이 충분히 쌓인 상황에서 이용자의 행동 이력을 셀프 어텐션(self-attention) 기법으로 분석해 이용자의 향후 반응을 예측한다. 두 번째 연구는 해당 게임에서의 기록이 없는 상황에서 여러 게임에 걸친 다운로드 이력과 협업 필터링(collaborative filtering)을 활용한다. 두 연구는 다루는 데이터의 상황과 사용 기법이 다르지만, 이용자의 행동 기록을 예측 대상에 맞는 모델과 신호로 어떻게 설계하느냐가 예측 성능을 좌우한다는 점을 공통으로 보여 준다.

뉴스 플러스

1. 서론 : 게임 이용자 행동 데이터와 AI 예측 기술

• 게임 이용자 행동 데이터 분석의 배경

디지털 게임 산업의 수익 구조는 과거의 일회성 패키지 판매에서 벗어나 부분 유료화(F2P: free-to-play)와 서비스형 게임(GaaS: games as a service) 모델을 중심으로 재편되고 있다. 이용자는 게임을 무료로 이용하는 대신 게임 내에서 아이템을 구매하거나 추가 혜택을 위한 구독료를 지불하며, 이러한 인게임 거래(in-game transaction)가 개발사 매출의 중심을 이루고 있다. 이에 따라 이용자의 결제 행동을 예측하는 일이 서비스 운영의 주요 과제 중 하나가 되었다.

게임 안에서는 소수의 이용자가 전체 매출의 대부분을 차지하며, 결제에 이르지 않은 이용자의 상당수는 게임을 떠나는 경향을 보인다.¹⁾ 따라서 어떤 이용자가 구매에 근접해 있는지, 또 어떤 게임에 얼마를 지출할지를 예측하는 일은 이용자 유지와 자원 배분의 중요한 근거가 된다. 이러한 예측의 바탕이 되는 것은 이용자가 게임과 상호작용하며 남기는 행동 기록이다. 여기에는 게임 접속 여부, 플레이 시간, 플레이 세션의 수, 통과한 레벨, 게임 내 재화의 획득과 소비 내역 등 게임 내부의 활동뿐 아니라, 어떤 게임을 내려받았는지와 같은 게임 외부의 기록도 포함된다. 이러한 데이터에는 이용자의 행동과 선호가 담겨 있으며, 이를 종합하면 개별 이용자나 집단 단위의 소비 패턴을 통계적으로 추론할 수 있다.

• 상반된 데이터 상황과 예측 과제

행동 데이터를 활용한 예측은 데이터가 얼마나 확보되어 있는지에 따라 서로 다른 기술적 과제를 낳는다. 첫 번째 경우는 이용자가 게임을 오래 이용하며 방대한 행동 이력을 축적한 상황이다. 이 경우 충분한 데이터를 확보할 수 있으나, 오래 축적된 이력 가운데 어느 시점의 행동이 미래의 결제와 연관되는지를 분별하기가 어렵다. 이용자의 구매 결정은 최근의 몇몇 활동만으로 정해지지 않으며, 오래전 시점의 행동이 최근 흐름과 얹혀 결제로 이어지기도 한다. 이처럼 긴 시계열(time-series) 데이터에서 시점상 멀리 떨어진 행동들 사이의 연관성을 포착하는 데에는 기존 분석 기법이 한계를 보인다.²⁾

두 번째 경우는 예측 대상 게임에 대한 이용자의 기록이 없는 상황이다. 예를 들어 이용자가 새로 내려받은 게임에서는 아직 결제나 플레이 이력이 쌓이지 않았기 때문에, 해당 게임에서의 향후 예상 지출을 추정할 근거가 존재하지 않는다. 이때는 대상 게임의 자료가 아니라 이용자가 다른 곳에서 남긴 기록을 활용해야 한다.³⁾ 결국 게임사는 데이터가 방대하여 그 안에서 유의미한 신호를 분별하기 어려운 상황과, 예측 대상에 대한 데이터 자체가 없는 상황이라는 두 갈래의 과제를 마주하게 된다.

• AI 예측 모델의 발전과 두 연구

이러한 어려움에 대응하기 위해 데이터 분석 기법은 꾸준히 발전해 왔다. 초기 분석은 접속자 수, 평균 플레이 시간, 이탈률처럼 전체 이용자 집단의 경향을 파악하는 데 초점을 두었다. 이러한 거시 지표는 게임의 전반적인 상태를 진단하는 데 유용했으나 개인화된 접근이나 정밀한 예측에는 한계가 있었다. 이후 개별 이용자의 행동 순서를 분석하기 위해 순환 신경망(RNN: recurrent neural network)*이나 장단기 메모리(LSTM: long short-term memory)**와 같은 딥러닝 모델이 도입되면서 시계열 분석에 진전이 있었으나, 시퀀스가 길어질수록 멀리 떨어진 정보 사이의 관계를 학습하는 데에는 여전히 한계를 보였다.⁴⁾

1) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

2) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

3) Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv 2024.04.12., <https://arxiv.org/abs/2404.08301>

4) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

본 보고서는 상반된 두 데이터 상황에 각각 대응하는 두 건의 연구를 살펴본다. 첫 번째 연구 (Kovačević 외, 2024)는 한 게임에 이용자의 행동 이력이 충분히 축적된 상황을 다룬다. 이 연구는 시퀀스 내 모든 요소 사이의 관계를 한꺼번에 고려해 중요한 부분에 더 큰 가중치를 두는 셀프 어텐션 (self-attention) 기법으로, 향후 3일, 5일, 7일 안에 구매가 일어날지를 예측한다. 두 번째 연구 (Sun 외, 2024)는 예측 대상 게임에 대한 기록이 없는 신규 상황을 다룬다. 이 연구는 이용자가 다른 게임에서 남긴 다운로드 이력을 단서로 취향이 비슷한 이용자를 찾아내는 협업 필터링(collaborative filtering) 기법으로 지출의 규모를 예측한다. 두 연구는 모두 이용자의 행동 기록을 분석하지만, 분석에 쓸 기록이 얼마나 확보되어 있는지는 서로 다르다. 이어지는 본문에서는 각각의 조건에서 두 연구가 어떻게 작동하며 어떤 성과와 한계를 보이는지를 기술적 관점에서 분석한다.

* 순환 신경망(RNN: recurrent neural network): 데이터를 순서대로 처리하며 앞선 정보를 다음 단계의 계산에 반영하는 신경망

** 장단기 메모리(LSTM: long short-term memory): 순환 신경망의 한 종류로, 어떤 정보를 오래 기억하고 어떤 정보를 잊을지 조절하는 장치를 두어 순환 신경망의 한계를 보완한 모델. 기존 순환 신경망이 시퀀스가 길어질수록 앞부분의 정보를 잃는 문제를 완화하도록 고안됨

II. 본론 1: 행동 이력이 축적된 상황에서의 예측

• 연구 개요와 대상 게임

첫 번째 연구는 한 게임에서 이용자의 행동 이력이 충분히 축적된 상황을 다룬다. 분석 대상은 모바일 퍼즐 게임 'Woka Woka'로, 2022년 10월부터 2023년 4월 사이에 가입한 이용자들의 행동을 가입 후 90일 동안 관찰한 기록이다. 관찰 대상으로 등록된 이용자는 977,287명이었으나, 이 가운데 한 번이라도 구매한 이력이 있는 이용자는 13,765명으로 전체의 1.4%에 그쳤다. 해당 연구는 구매 이력이 없는 이용자를 분석 대상에서 제외하고, 구매 경험이 있는 이용자를 대상으로 첫 구매 시점부터 마지막 구매 이후 7일까지의 행동을 추적하였다. 관찰된 데이터에는 하루하루의 접속 여부와 구매 여부, 플레이 시간, 세션 수, 통과한 레벨 수, 게임 내 재화의 구매·획득·소비량과 잔여 보유량 등이 포함된다. 동 연구는 이러한 행동 이력을 바탕으로 이용자가 향후 3일, 5일, 7일 안에 구매할지를 예측하고자 한다.⁵⁾

여기서는 행동의 종류만큼이나 그 행동이 나타나는 순서가 중요하다. 이용자의 행동은 접속, 플레이, 아이템 획득, 상점 방문으로 이어지는 시간적 순서 속에서 읽을 때 의미 있는 패턴으로 해석되기 때문이다. 개별 행동을 따로 떼어 파악하는 대신 시간 순으로 이어진 행동의 전체 맥락을 학습하려면, 이러한 순차 데이터(sequential data)를 다루는 데 적합한 방법이 필요하다. 이러한 이유에서 동 연구는 트랜스포머 (transformer)* 신경망과 그 핵심 기법인 셀프 어텐션 (self-attention)을 채택하였다.

* 트랜스포머(transformer): 셀프 어텐션 기법을 핵심으로 하는 신경망 모델. 순서대로 이어지는 데이터를 다룰 때 어느 부분이 서로 관련이 깊은지를 찾아내는 데 강점이 있으며, 시간 간격이 멀리 떨어진 정보 사이의 관계도 잡아낼 수 있어 언어 번역을 비롯한 다양한 분야에 활용됨

5) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

• 행동 이력의 토큰화

트랜스포머 모델이 이용자의 행동 이력을 학습하려면, 먼저 그 이력을 모델이 처리할 수 있는 형태로 바꾸어야 한다. 이 과정을 토큰화(tokenization)라 하며, 연속적이고 다양한 값을 갖는 데이터를 한정된 종류의 표준화된 기호, 곧 토큰(token)으로 변환하는 것을 말한다. 언어 모델이 문장을 단어 단위의 토큰으로 나누어 처리하듯, 이용자의 행동 이력도 토큰의 나열로 바꾸면 같은 방식으로 학습할 수 있다.

[표 1] 이용자의 하루를 기록하는 13개 특성

	특징	설명
게임 참여도	활동 여부	해당 일에 게임 활동이 있었는지 여부(0/1)
	누적 플레이 시간	가입 시점부터 해당 일까지의 총 플레이 시간
	일일 세션 수	해당 일에 발생한 게임 접속 세션의 수
플레이 숙련도	누적 통과 레벨 수	가입 이후 통과한 레벨의 총수
	레벨 통과 비율	해당 일 플레이한 서로 다른 레벨 가운데 통과한 레벨의 비율
	레벨 재도전 비율	해당 일 플레이한 레벨 가운데 이전에 플레이했던 레벨의 비율
	되감기 사용 횟수	해당 일 사용한 되감기 기능의 횟수
	부스터 사용 횟수	해당 일 사용한 부스터 아이템의 횟수
과금 행동	구매 여부	해당 일에 결제가 발생했는지 여부(0/1), 모델의 예측 목표
	재화 구매량	해당 일 현금 결제로 구매한 게임 내 재화의 양
	재화 획득량	해당 일 플레이 보상으로 획득한 재화의 양
	재화 사용량	해당 일 소비한 재화의 양
	재화 보유량	해당 일 종료 시점에 남아 있는 재화의 양

출처: Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, vol. 12, 2024.10.10, <https://ieeexplore.ieee.org/document/10713356> (논문 정보 기반 재구성)

해당 연구는 이용자의 하루를 13개 특성(feature) 값으로 기록하고, 각 특성값을 하나의 토큰으로 변환하였다(표 1). 접속 여부나 구매 여부처럼 값이 0 또는 1뿐인 특성은 두 종류의 토큰으로 표현되고, 플레이 시간이나 재화량처럼 연속적인 수치를 갖는 특성은 그 값을 그대로 쓰지 않고 네 종류의 토큰 중 하나로 표현된다.

연속적인 값을 몇 개의 구간으로 나누어 한정된 종류의 값으로 바꾸는 과정을 이산화(discretization)라 한다. 예컨대 플레이 시간은 네 구간으로 구분되어 각각 하나의 토큰에 대응된다. 그날 게임을 하지 않아 값이 없는 경우, 다른 값들과 동떨어질 만큼 유난히 큰 값(이상치)인 경우, 그리고 그 외 정상적인 값을 낮은 구간과 높은 구간으로 나눈 두 경우가 이에 해당한다. 어느 값까지를 낮은 구간으로 보고 어디부터 높은 구간으로 볼지, 그 경계를 정하는 데에는 K-평균 군집화(K-means clustering)*와 같은 방법이 활용된다. 다만 K-평균 군집화는 이상치에 민감하므로, 먼저 사분위 범위(IQR: interquartile range)** 방법으로 이상치를 분리한 뒤 나머지 정상값의 구간 경계를 정한다. 이 과정을 거치면 연속적이고 제각각이던 행동 기록이 한정된 종류의 기호로 정리된다. 이때 세부적인 수치 정보는 일부 사라지지만, 그 대신 모델은 행동 패턴의 구조를 학습하기에 적합한 형태의 데이터를 얻는다.

토큰화를 거치면 이용자의 하루는 13개 토큰의 묶음이 되고, 가입 시점부터 이어지는 전체 이력은 이 하루 단위 묶음을 시간 순으로 이어 붙인 하나의 긴 시퀀스(sequence)가 된다. 플레이 기간은 이용자마다 다르므로 시퀀스의 길이도 서로 다르다. 이처럼 길이가 제각각인 이력을 하나의 고정된 길이로 맞추지 않고도 유연하게 처리할 수 있다는 점이 트랜스포머 모델의 주요 강점이다.⁶⁾

* K-평균 군집화(K-means clustering): 데이터를 서로 가까운 값끼리 K개의 묶음으로 자동 분류하는 기계학습 기법. 여기서는 각 특성의 수치 범위를 비슷한 값끼리 구간으로 나누는 데 쓰임

** 사분위 범위(IQR: interquartile range): 데이터를 크기순으로 늘어놓았을 때 가운데 절반이 퍼져 있는 정도를 나타내는 값. 이 범위에서 크게 벗어난 값을 이상치로 판별하는 데 쓰임

• 셀프 어텐션 메커니즘

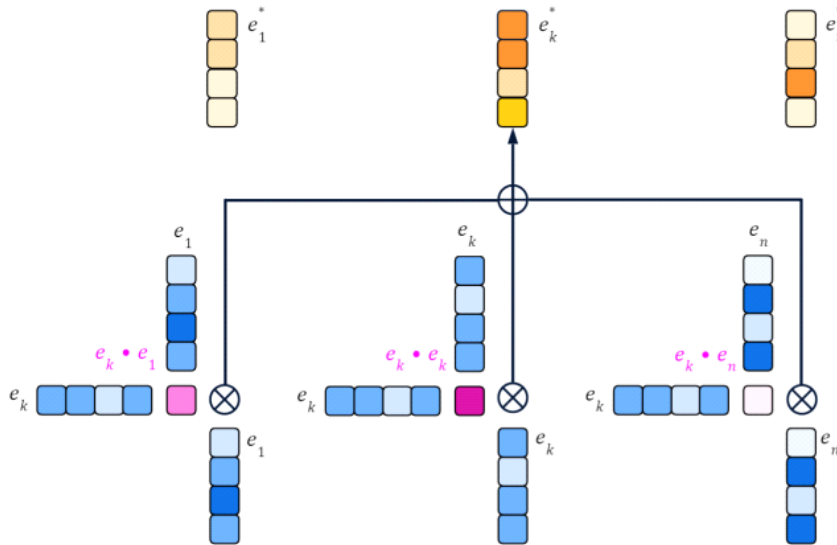
이용자의 행동 이력을 토큰 시퀀스로 변환하더라도, 장기간에 걸친 기록 가운데 어떤 부분이 미래의 구매 결정과 관련이 깊은지를 가려내야 정확한 예측이 가능하다. 이를 위해 활용되는 것이 셀프 어텐션 메커니즘이다. 셀프 어텐션이란 모델이 하나의 시퀀스를 분석할 때 그 안의 개별 토큰들 사이의 연관성을 학습하여, 어떤 요소에 더 주의를 기울일지 정하는 방식이다. 각 토큰을 의미 공간의 벡터로 표현한 뒤 모든 토큰 쌍 사이의 유사도를 동시에 계산하고, 그 결과를 바탕으로 각 토큰을 시퀀스 전체의 맥락이 반영된 새로운 벡터로 변환한다. 예컨대 기준이 되는 한 토큰(e_k)은 시퀀스 내 모든 토큰과의 유사도를 반영하여 맥락이 담긴 새로운 표현(e_k^*)으로 바뀐다(그림 1).⁷⁾

이 과정에서 셀프 어텐션은 모든 행동 기록에 같은 무게를 두지 않는다. 예를 들어 보유한 게임 재화가 바닥난 날과 그 직후 플레이가 급격히 늘어난 며칠은, 평이한 플레이가 이어진 수십 일보다 향후 구매를 예측하는 데 더 중요한 단서가 될 수 있다. 셀프 어텐션은 이처럼 단서가 되는 토큰의 조합을 데이터로부터 찾아내고, 그 관계를 예측에 반영한다.

6) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

7) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

[그림 1] 셀프 어텐션 메커니즘의 작동 원리

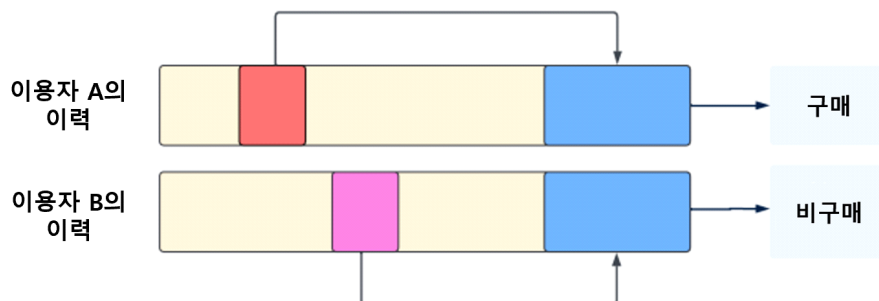


출처: Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, vol. 12, 2024.10.10, <https://ieeexplore.ieee.org/document/10713356> (논문 정보 기반 재구성)

이러한 방식 덕분에 셀프 어텐션은 긴 행동 이력 속에서 시점 간 거리가 먼 행동들 사이의 관계까지 포착할 수 있다. 해당 연구는 최근의 행동 패턴이 비슷한 두 이용자라도 그 이전의 전체 이력이 다르면 구매 결정이 다르게 나타날 수 있다고 본다(그림 2). 직전 며칠의 행동뿐 아니라 몇 주, 몇 달 전의 행동까지 함께 고려할 때 예측이 정확해지며, 셀프 어텐션은 이러한 장거리 의존성(long-range dependency)*에 적합한 기법이다.⁸⁾

* 장거리 의존성(long-range dependency): 순서가 있는 데이터에서 서로 멀리 떨어진 요소가 의미상 밀접하게 연관되는 성질. 여기서는 며칠 전이 아니라 몇 주·몇 달 전의 행동이 현재의 구매 결정과 관련되는 경우를 가리킴

[그림 2] 과거 이력에 따라 달라지는 구매 결정



출처: Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, vol. 12, 2024.10.10, <https://ieeexplore.ieee.org/document/10713356> (논문 정보 기반 재구성)

8) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

• 예측 성능과 불균형 데이터에서의 평가

학습을 마친 모델은 특정 이용자의 행동 이력을 입력받아 향후 3일, 5일, 7일 안에 구매가 발생할지를 예측한다. 해당 연구는 동일한 데이터에 랜덤 포레스트(random forest)*, XGBoost**, 다층 퍼셉트론(MLP: multilayer perceptron)*** 등 널리 쓰이는 기계학습 모델을 함께 학습시켜 성능을 비교하였다. 그 결과 셀프 어텐션 기반 모델이 세 예측 기간 모두에서 가장 높은 성능을 기록하였다. 7일 예측을 기준으로 보면 F1 점수가 0.8495로 차순위인 XGBoost(0.8298)를 앞섰고, AUROC(area under the ROC curve)와 매슈스 상관계수(MCC: Matthews correlation coefficient)에서도 가장 높은 값을 보였다(표 2). 5일과 3일 예측에서도 셀프 어텐션 기반 모델의 우위는 동일하게 나타났다.

* 랜덤 포레스트(random forest): 여러 개의 의사결정나무를 만들어 그 예측을 종합하는 기계학습 기법. 단일 모델보다 안정적이고 정확한 예측을 얻기 위해 널리 쓰임

** XGBoost(extreme gradient boosting): 약한 예측 모델을 순차적으로 결합해 성능을 끌어올리는 부스팅 계열의 기계학습 기법. 정형 데이터 예측에서 높은 정확도로 널리 활용됨

*** 다층 퍼셉트론(MLP: multilayer perceptron): 여러 층으로 이루어진 기본적인 인공신경망 구조. 입력과 출력 사이에 중간층을 두어 복잡한 비선형 관계를 학습함

[표 2] 구매 예측 모델의 성능 비교*

모델	F1	AUROC	MCC
랜덤 포레스트	0.8059	0.9028	0.6121
XGBoost	0.8298	0.9231	0.6600
다층 퍼셉트론	0.7776	0.8875	0.5659
트랜스포머(셀프 어텐션)	0.8495	0.9348	0.6989

* 성능 지표(F1 점수·AUROC·MCC): 분류 모델의 성능을 수치로 평가하는 지표. F1 점수는 실제 구매자를 빠짐없이 찾아내는가와 구매자로 예측한 것이 실제로 맞는가를 함께 반영함. AUROC는 구매자와 비구매자를 구별하는 능력을 나타냄. MCC는 맞힌 예측뿐 아니라 놓치거나 잘못 분류한 예측까지 함께 반영하여 분류 성능을 -1에서 +1 사이 값으로 나타냄. 한쪽으로 크게 치우친 데이터에서도 예측의 정확성을 한쪽에 쏠리지 않게 평가하는 데 쓰임

출처: Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, vol. 12, 2024.10.10, <https://ieeexplore.ieee.org/document/10713356> (논문 정보 기반 재구성)

이 연구가 F1 점수와 AUROC와 함께 MCC를 평가 지표로 삼은 것은 데이터의 불균형 때문이다. 앞서 살펴보았듯 구매 경험자는 전체 이용자의 1.4%에 그쳐, 구매와 비구매의 비율이 크게 치우쳐 있다. 이처럼 한쪽으로 치우친 데이터에서는 단순 적중률이 실제 예측 능력을 왜곡하여 보여줄 수 있다. MCC는 맞힌 예측뿐 아니라 놓치거나 잘못 분류한 예측까지 네 가지 경우를 모두 반영하여 성능을 -1에서 +1 사이의 값으로 나타내므로, 한쪽으로 치우친 데이터에서도 특정 부류에 쏠리지 않고 성능을 평가하는 데 적합하다. 셀프 어텐션 기반 모델이 MCC에서 우위를 보였다는 것은 다수를 이루는 부류뿐 아니라 소수인 부류까지

비교적 정확히 가려냈음을 의미한다. 이러한 우위는 토큰화 과정에서 세부 수치가 일부 손실되었음에도 나타난 결과로, 개별 수치의 정밀함보다 행동 이력 전체의 장기 패턴을 포착하는 능력이 구매 예측에서 더 크게 작용함을 보여 준다.⁹⁾

• 모델의 기술적 한계

셀프 어텐션 기반 예측 모델은 높은 정확도를 보이지만, 몇 가지 기술적 한계도 함께 지닌다. 우선 연구진이 스스로 밝힌 한계로, 모델의 성능이 토큰화 방식에 민감하다는 점을 들 수 있다. 연속적인 행동 데이터를 몇 개의 구간으로 나누어 토큰으로 바꾸는 과정에서 구간을 어떻게 설정하느냐에 따라 결과가 달라질 수 있으며, 연구진은 더 나은 토큰화 방법의 탐색을 후속 연구 과제로 제시하였다. 이는 한정된 종류의 기호, 곧 이산화된 표현으로도 장기 패턴을 포착한다는 이 모델의 강점이 역으로 그 기호를 어떻게 설계하느냐에 그만큼 좌우된다는 뜻이기도 하다.

인공지능 모델 일반의 한계로 지적되는 블랙박스(black box) 특성도 있다. 셀프 어텐션 메커니즘은 모델이 어떤 행동 토큰에 더 주목했는지에 대한 단서는 제공하지만, 수많은 행동 조합 가운데 왜 특정 패턴이 구매의 핵심 신호로 학습되었는지, 그 인과관계까지 설명해 주지는 않는다. 또한 모델의 성능은 학습 데이터의 품질과 구성에 좌우된다. 대규모 업데이트나 새로운 이벤트처럼 과거 데이터에 없던 콘텐츠가 등장하면 이용자의 행동 패턴이 단기간에 변하면서 기존 모델의 예측 정확도가 일시적으로 낮아질 수 있다. 실제 운영 환경에서 이러한 모델이 최신 데이터를 반영한 주기적 재학습을 거쳐 운용되는 경우가 많은 것도 이러한 이유 때문이다.¹⁰⁾

II. 본론 2: 행동 이력이 없는 상황에서의 예측

• 콜드 스타트 : 기록이 없는 상황에서의 예측

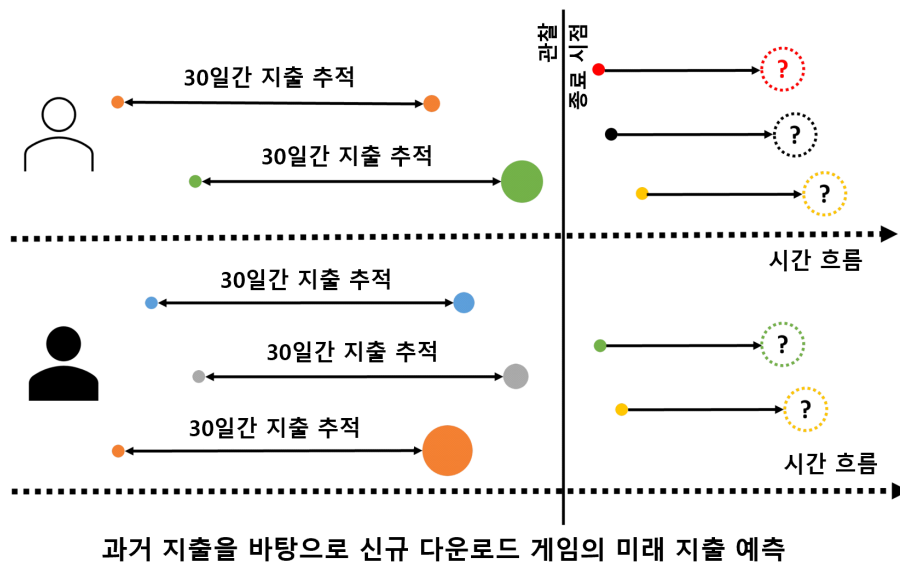
두 번째 연구가 다루는 상황은 본론 1과 크게 다르다. 본론 1의 모델이 한 게임에 이용자의 행동 이력이 충분히 축적된 상황을 전제했다면, 이 연구는 그러한 이력 자체가 없는 경우를 다룬다. 이용자가 새로 내려받은 게임에서 앞으로 얼마를 지출할지를 그 게임에서의 과거 기록이 없는 상태에서 추정해야 하기 때문이다. 이처럼 기록이 없는 신규 이용자나 신규 항목에 대해 예측을 수행하는 문제를 추천·예측 분야에서는 콜드 스타트(cold start)*라 부른다. 게임사가 예측을 가장 필요로 하는 시점이 이용자가 게임을 막 내려받은 초기이므로 이 문제는 실무에서 중요하게 다루어진다.

9) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

10) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>

본론 1의 모델이 한 게임 안에 축적된 행동 이력을 읽어냈다면, 본론 2의 연구가 다루는 상황은 대상 게임에 대한 이력 자체가 확보되지 않은 경우에 해당한다. 해당 게임에서의 기록은 없지만 이용자가 이전에 다른 게임들에서 남긴 이력은 활용할 수 있으며, 이를 단서로 새로 받은 게임에서의 지출을 예측한다(그림 3).¹¹⁾

[그림 3] 과거 지출 이력에 기반한 미래 지출 예측



출처: Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12, <https://arxiv.org/abs/2404.08301>

• 이용자 ID 없이 작동하는 협업 필터링

이처럼 참고할 이력이 없는 상황에서 예측을 수행하는 데 활용되는 것이 협업 필터링(collaborative filtering)이다. 협업 필터링이란 수많은 이용자의 행동 패턴에서 유사성을 파악하여 특정 이용자의 알려지지 않은 취향이나 관심사를 예측하는 기법이다. 취향이 비슷한 이용자들이 공통으로 선호한 항목이라면 해당 이용자도 선호할 가능성이 높다고 추론하는 원리다. 온라인 쇼핑몰이 보여 주는 '이 상품을 함께 본 다른 고객의 추천 목록'이나, 콘텐츠 플랫폼이 비슷한 시청 이력의 이용자에게 권하는 작품이 대표적인 사례다.

다만 협업 필터링은 전통적으로 개별 이용자를 식별하는 ID를 기반으로 작동해 왔다. 이용자 간에 취향이 비슷한지를 따지려면 우선 각 이용자를 구별할 수 있어야 하기 때문이다. 해당 연구는 ID를 사용하는 대신 이용자가 과거에 내려받은 게임 앱의 목록을 취향을 대변하는 지표로 활용하는 방식을 제안하였다. 다운로드 이력을 신경망에 입력해 이용자별 선호를 나타내는 벡터 표현을 만들고, 이를 예측 대상 게임의 표현과 결합해 이용자와 게임 사이의 적합도, 곧 협업 신호(collaborative signal)를 산출하는 구조다.

11) Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12., <https://arxiv.org/abs/2404.08301>

어떤 게임들을 내려받아 왔는지가 그 이용자의 취향을 보여 주므로, 식별 정보 없이도 취향이 비슷한 이용자들의 소비 패턴이 예측에 반영된다. 이러한 설계는 개인정보 보호에 기여할 뿐 아니라, ID 목록을 따로 유지할 필요가 없어 실시간으로 유입되는 데이터로 모델을 갱신하는 온라인 학습 환경에도 적합하다.¹²⁾

• 지출 데이터의 불확실성과 레이블 표준화

지출 규모의 예측이 까다로운 이유는 데이터의 분포에 있다. 해당 연구가 한 달간 수집한 약 2,970만 건의 이용자-게임 기록 가운데 실제 지출이 발생한 표본은 약 62만 건으로, 전체의 2%를 조금 넘는 수준이었다. 게다가 지출이 있는 표본 안에서도 편차가 매우 컸다. 평균 지출액이 227(단위 비공개)인 반면 표준편차는 2,130에 이르고 최댓값은 42만을 넘어, 소수의 고액 지출이 전체 분포를 크게 좌우하는 형태였다. 이러한 값을 가공 없이 학습 목표로 삼으면 극단값이 학습 과정을 지배하여 모델이 안정적인 패턴을 익히기 어렵다.

해당 연구는 이 문제를 레이블 표준화(label standardization)로 다루었다. 지출액의 절대 크기가 아니라 그 지출이 이용자와 게임 각각의 통상적인 수준에 비추어 얼마나 큰지를 학습 목표로 삼는 방식이다. 구체적으로는 같은 지출액을 두 가지 기준으로 나누어 평가한다. 게임 측면에서는 해당 게임에서 그동안 발생한 지출 분포에 비추어 그 지출이 얼마나 큰지를 가늠하고, 이용자 측면에서는 그 이용자가 과거 180일간 결제해 온 평균 금액에 비추어 가늠한다. 이렇게 얻은 두 값을 각각 절반씩 반영해 그 평균을 최종 학습 목표로 삼는다. 같은 금액이라도 소액 결제 위주의 이용자에게는 큰 지출이지만, 고액 결제가 흔한 게임에서는 평범한 지출에 해당한다. 이러한 상대적 차이가 표준화를 거쳐 학습 목표에 반영된다. 실험에서도 이용자와 게임 양측을 함께 고려한 표준화가 한쪽만 적용하거나 원래 값을 그대로 쓰는 경우보다 일관되게 나은 성능을 보였다.¹³⁾

• 예측 성능 : 오프라인 실험과 실제 운영

해당 연구는 이 모델을 두 단계에 걸쳐 검증하였다. 우선 오프라인 실험에서 XGBoost, 구글의 ZILN, 콰이쇼우의 MDME* 등 기존 기법과 비교한 결과, 결정계수(R^2) 기준으로 기존 운영 모델보다 17.11% 높은 값을 기록하며 비교 대상 가운데 가장 우수하였다(표 3). 이러한 예측 정확도의 개선은 실제 운영 지표로도 확인되었다. 연구진은 화웨이 게임스토어(Huawei Game Store)의 게임 추천 영역에서 이용자를 두 집단으로 나누어, 한쪽 집단에만 이 모델을 적용하는 A/B 테스트**를 약 2주간 수행하였다. 그 결과 제안 모델을 적용한 집단에서 신규 게임 관련 매출이 기존 모델을 적용한 집단보다 50.65% 높게 나타났다.

* ZILN·MDME: 게임·앱 이용자의 생애가치(LTV: lifetime value), 즉 한 이용자가 서비스를 이용하는 동안 발생시키는 총지출을 예측하기 위해 개발된 딥러닝 모델. 각각 구글과 콰이쇼우가 제안함

** A/B 테스트: 이용자를 두 집단으로 나누어 한쪽에만 새로운 방식을 적용하고 다른 쪽에는 기존 방식을 유지한 뒤, 두 집단의 반응을 비교해 효과를 확인하는 검증 방법. 두 집단의 차이가 적용한 방식의 차이에서 비롯되도록 설계함

12) Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12., <https://arxiv.org/abs/2404.08301>

13) Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12., <https://arxiv.org/abs/2404.08301>

다만 같은 기간 게임 다운로드 수가 오히려 0.88% 감소했다. 연구진은 이를 모델이 이용자가 실제로 만족할 만한 게임을 추천한 결과로 해석하였다. 내려받기만 하고 마는 게임이 줄어든 대신, 추천을 통해 받은 게임이 지출로 이어졌다는 의미다. 이러한 결과는 식별 정보 없이 다운로드 이력만으로도 이용자의 취향을 예측할 수 있음을 보여 준다.¹⁴⁾

[표 3] 지출 예측 모델의 오프라인 성능 비교***

모델	RMSE	R ²	AUC
XGBoost	246.40	0.0257	0.8336
ZILN(구글)	242.56	0.0558	0.8601
MDME(콰이쇼우)	243.71	0.0440	0.8482
기존 운영 모델	240.16	0.0744	0.8629
제안 모델(협업 신호 결합)	238.50	0.0871	0.8693

* 성능 지표(RMSE·R²·AUC): 예측 모델의 성능을 수치로 평가하는 지표. RMSE는 예측한 지출액이 실제와 얼마나 어긋나는지를 나타내며 낮을수록 우수함. R²(결정계수)는 모델이 실제 지출의 변동을 설명하는 정도를 나타내며, 값이 높을수록 우수함. AUC는 구매자와 비구매자를 구별하는 능력을 나타내며, 값이 높을수록 우수함

출처: Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12, <https://arxiv.org/abs/2404.08301> (논문 정보 기반 재구성)

• 비식별 설계의 의의와 데이터 편향

이 모델의 두드러진 특징은 개인정보 보호를 사후 조치가 아니라 설계 단계에서부터 반영하였다는 점이다. 개별 이용자를 식별하는 ID를 사용하지 않고 다운로드 이력이라는 행동 기록만으로 선호를 표현하므로, 예측 과정에서 특정 개인을 추적할 필요가 없다.¹⁵⁾ 이용자 행동 데이터의 분석은 개인정보 침해 우려와 자주 부딪혀 왔는데, 비식별 상태에서도 충분한 예측력을 확보할 수 있다는 점은 이러한 우려를 덜면서 기술의 적용 범위를 넓히는 의미 있는 결과다. 다만 다운로드 이력은 어디까지나 취향을 간접적으로 보여 주는 지표라는 한계를 함께 지닌다. 무료 게임을 가볍게 받아 보는 행동이 실제 선호와 항상 일치하지는 않으며, 노출이 많은 인기 게임일수록 다운로드 이력에 자주 등장하므로 모델이 이용자의 고유한 취향보다 인기 게임의 영향을 반영하도록 학습될 가능성도 있다.

14) Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12., <https://arxiv.org/abs/2404.08301>

15) Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12., <https://arxiv.org/abs/2404.08301>

데이터 편향 문제도 같은 맥락에 있다. 학습 데이터가 특정 장르나 특정 소비 성향의 이용자에게 치우치면 모델이 형성하는 예측의 기준 자체가 왜곡될 수 있다. 실제 운영에서 편향의 측정과 보정이 모델 관리의 한 부분으로 다루어지는 것도 이 때문이다.

III. 결론 및 전망

• 데이터 상황과 예측 방법의 대응

본 보고서는 게임 이용자의 행동 데이터를 분석하는 AI 예측 기술을 두 건의 연구를 통해 살펴보았다. 두 연구는 이용자의 행동 기록을 분석한다는 점에서는 같으나, 다루는 상황과 예측의 성격은 서로 달랐다. 첫 번째 연구는 한 게임에 이용자의 행동 이력이 충분히 축적된 상황에서 그 이력을 셀프 어텐션 기법으로 분석해 향후 며칠 안에 구매가 일어날지를 파악하였다. 두 번째 연구는 이용자가 새로 내려받은 게임처럼 해당 게임에서의 기록이 없는 콜드 스타트 상황에서 여러 게임에 걸친 다운로드 이력과 협업 필터링을 활용해 지출 규모를 추정하였다. 이처럼 두 연구는 서로 다른 상황을 각기 다른 방법으로 다루면서도, 모두 기존 기법을 웃도는 예측 성능을 보였다. 특히 두 번째 연구는 실제 서비스 운영 환경에서도 성능을 검증하였다.

• 데이터 가공의 중요성과 그 이면

본 보고서가 다룬 두 연구를 함께 살펴보면 한 가지 공통된 원리가 확인된다. 두 연구 모두 원시 데이터를 모델이 학습하기 좋은 형태로 다듬는 단계에 주의를 기울였다는 점이다. 첫 번째 연구는 연속적이고 값이 다양한 행동 기록을 토큰화를 거쳐 한정된 종류의 기호로 변환하였고, 두 번째 연구는 편차가 큰 지출액을 그대로 사용하는 대신 이용자와 게임 각각의 통상적 수준에 비추어 그 크기를 나타내는 값으로 표준화하였다. 다루는 문제도 방법도 서로 달랐으나, 두 연구 모두 데이터의 고유한 특성에 대응하는 처리를 모델 설계의 중요한 전제로 삼았다는 점에서 공통된다.

이 공통점은 곧 공통의 한계이기도 하다. 첫 번째 연구는 토큰화 방식에 두 번째 연구는 표준화 기준에 성능이 민감하게 반응하였다. 데이터를 다듬는 설계가 성능을 높이는 만큼 그 설계가 어긋나면 성능도 함께 저하될 수 있다. 두 모델은 예측의 근거를 충분히 설명하기 어려운 블랙박스 특성과 학습 데이터의 구성에 따라 예측 기준이 치우칠 수 있는 편향 문제도 함께 지닌다.

이에 따라 행동 데이터 예측 모델은 실제 운영 환경에서 지속적인 점검과 보완을 전제로 운용된다. 모델이 어떤 행동 패턴을 근거로 특정 예측에 이르렀는지를 사람이 이해할 수 있는 형태로 제시하려는 설명 가능 AI(XAI: explainable AI) 연구도 이러한 한계를 좁히는 방향으로 이어지고 있다. 예측의 근거가 함께 제시될수록 예측 결과에 대한 신뢰가 높아지고 그 활용 범위도 넓어질 것으로 전망된다.

• 데이터 가공과 학습 방식의 발전 과제

두 연구의 성능이 데이터를 다루는 방식에 달려 있다는 사실은 앞으로의 발전 여지 또한 그 지점에 있음을 보여 준다. 첫 번째 연구진은 연속적인 행동 값을 토큰으로 변환하는 과정이 성능을 좌우하는 요인임을 확인하고, 더 나은 토큰화 방법의 탐색을 후속 과제로 제시하였다.¹⁶⁾ 모델 구조의 개선과 더불어 행동 기록을 어떤 기호로 변환하는지를 정교화하는 일 또한 예측 정확도를 높이는 요소로 볼 수 있다.

학습에 활용하는 데이터의 범위에서도 발전의 여지가 있다. 두 번째 연구가 개인을 식별하지 않고도 지출 규모를 예측한 것처럼, 학습에 어떤 데이터를 포함하고 제외할지를 정하는 설계는 성능뿐 아니라 이용자 보호와도 연관된다. 또한 데이터 의존성과 편향의 문제는 학습 데이터를 다양하게 확보하고 시간에 따른 행동 변화를 반영함으로써 완화될 수 있다. 이용자가 게임과 상호작용하며 남기는 행동 기록이 꾸준히 축적되는 만큼, 그 기록을 어떤 형태로 다듬어 무엇을 예측할 것인지에 관한 연구도 이어질 것으로 보인다. 게임 이용자의 행동 데이터를 활용한 예측 기술은 데이터를 다루는 방식이 정교해질수록 그 정확도와 적용 범위를 함께 넓혀 갈 것으로 전망된다.

참고문헌

- Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>
- Peijie Sun 외 7인, "Collaborative-Enhanced Prediction of Spending on Newly Downloaded Mobile Games under Consumption Uncertainty", arXiv, 2024.04.12., <https://arxiv.org/abs/2404.08301>

16) Miloš A. Kovačević 외 3인, "Predictive Analytics of In-Game Transactions: Tokenized Player History and Self-Attention Techniques", IEEE Access, 2024.10.10., <https://ieeexplore.ieee.org/document/10713356>