

추천 기준이 바뀌면 같은 입력도 다른 결과가 나올까?

똑똑 · 에디터 나루 · 발행 2026.09.28

입력·규칙·목표를 나누면 추천의 결과가 보입니다.

개념 설명과 자체 제작한 적용 사례를 함께 읽습니다. 가상 사례·수치는 해당 위치에 표시했습니다.

쉬운 설명

급식 줄을 키순으로 세울지, 먼저 온 순서로 세울지에 따라 맨 앞 사람이 달라집니다. 둘 다 정해진 순서대로 실행할 수 있지만, 어떤 기준이 좋은지는 다른 질문입니다.

핵심

- 알고리즘은 문제를 해결하는 절차로 이해할 수 있습니다.
- 추천 시스템의 목표와 데이터는 결과에 영향을 줍니다.
- 절차의 일관성과 선택 기준의 타당성은 구분해야 합니다.

정의 · 절차를 정하는 일, 목표를 고르는 일

알고리즘은 문제 해결을 위한 절차입니다. 책을 제목순으로 정렬하거나 여러 길 중 짧은 길을 찾는 일에도 알고리즘을 생각할 수 있습니다. AI가 등장하면서 생긴 개념은 아니며, 모든 알고리즘이 학습을 통해 만들어지는 것도 아닙니다.

이 글은 같은 입력을 서로 다른 목표와 규칙으로 처리할 때의 차이를 비교합니다. 입력과 규칙이 모두 같은 결정적 알고리즘은 같은 결과를 내며, 무작위 선택을 사용하는 절차는 별도로 구분해야 합니다.

먼저 문제를 어떻게 정의했는지가 중요합니다. “학생에게 좋은 글을 추천한다”는 목표는 그 자체로 계산 규칙이 되기 어렵습니다. 많이 클릭할 글인지, 이해를 돕는 글인지, 다양한 관점을 보여주는 글인지에 따라 수집할 데이터와 결과의 평가 방법이 달라집니다.

작동 방식 · 추천 결과 앞의 세 가지 선택

- 입력: 어떤 행동·자료를 관찰하는가
- 목표·규칙: 무엇을 좋은 결과로 세는가
- 평가: 누구에게 어떤 결과가 생겼는가

특정 서비스의 내부 설계가 아닌 추천을 읽기 위한 개념도입니다.

가상 사례 · 클릭만 잘하는 추천과 공부를 돕는 추천

가상의 독서 앱이 클릭 수만 높이는 제목을 추천한다고 해봅시다. 짧고 자극적인 글이 많이 늘리면 이런 글을 계속 선택할 수 있습니다. 반면 읽은 뒤 낯선 개념을 설명할 수 있는지를 목표로 삼으면 다른 자료가 필요합니다. 클릭이라는 관찰값이 이해라는 목표를 완벽히 대신하지 못하기 때문입니다.

학생이 어려운 글을 오래 열어두었다고 해서 깊이 읽었다고 단정할 수도 없습니다. 잠시 자리를 비웠거나 이해가 되지 않아 멈췄을 수 있습니다. 로그는 행동의 일부를 보여줄 뿐, 학습자의 생각을 직접 읽는 장치가 아닙니다.

구분 · 기계가 했다는 설명 다음에 물을 것

확인할 것	예시 질문
절차	같은 입력을 어떻게 처리하는가?
데이터	누구의 경험이 빠져 있는가?
목표	클릭·시간·이해 중 무엇을 최적화하는가?
이의 제기	결과를 설명받거나 바꿀 방법이 있는가?

기술을 평가하는 자체 질문이며 제품 인증 항목은 아닙니다.

비판적 읽기 · 오류가 없는 계산도 나쁜 선택을 할 수 있습니다

정해진 규칙을 틀리지 않고 실행해도 기준 자체가 부적절하면 결과는 문제가 될 수 있습니다. 과거 데이터에 기회 차이가 담겨 있는데 그대로 미래의 자격을 판단한다면 기존의 불리함이 반복될 가능성을 살펴야 합니다. 따라서 정확도 하나가 모든 품질을 대신하지 않습니다.

NIST의 생성형 AI 위험 자료는 그럴듯한 오류, 편향 등 여러 위험을 함께 다룹니다. 이것이 모든 시스템이 같은 방식으로 실패한다는 뜻은 아닙니다. 추천 알고리즘과 생성형 언어 모델의 작동을 동일하게 설명하지 않고, 실제 시스템이 하는 일에 맞춰 검증 질문을 골라야 합니다.

직접 생각해 보기

반 친구에게 책을 추천하는 규칙 두 가지를 만들고 각각 놓치는 점을 써보세요.

확인 설명

대출이 많은 책만 추천하면 접근은 쉽지만 새로운 관심사를 놓칠 수 있습니다. 관심 주제만 맞추면 익숙한 관점에 머물 수 있습니다. 규칙의 목적과 한계를 함께 쓰는 것이 핵심입니다.

알고리즘은 곧 AI인가요?

아닙니다. 정렬이나 계산 같은 절차도 알고리즘입니다. AI는 더 넓은 기술 맥락에서 여러 알고리즘과 데이터를 사용합니다.

추천 결과가 마음에 들지 않으면 오류인가요?

오류일 수도 있지만 시스템의 목표가 사용자의 기대와 다를 수도 있습니다. 어떤 기준으로 추천했는지 먼저 확인해야 합니다.

원문과 출처

- NIST · 알고리즘 사전 — <https://xlinux.nist.gov/dads/HTML/algorithm.html>: 알고리즘의 정의와 문제 해결 절차. 특정 추천 서비스의 내부 작동을 공개한 자료는 아닙니다.
- NIST · 생성형 AI 위험 관리 프로파일 (2024) — <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>: AI 600-1의 confabulation 등 위험 분류. 제품별 오류율을 제시하지 않습니다.

제작 방식과 정정 안내: <https://www.dokdok.co/about#editorial>

온라인 원문 · 수정된 내용은 아래 페이지에서 확인하세요.

<https://www.dokdok.co/reading/algorithm>