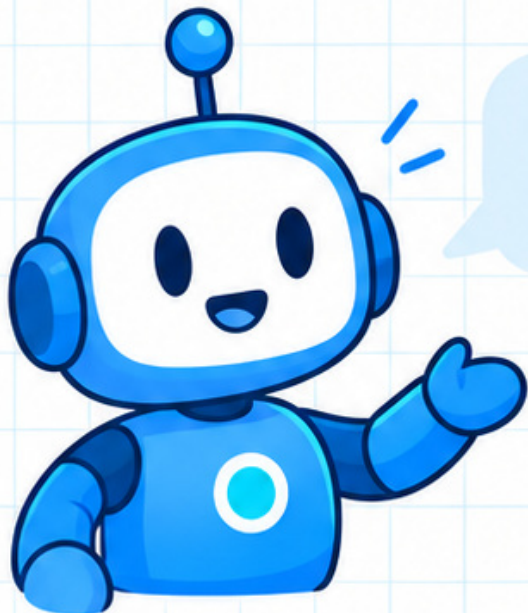


FIND CANDIDATES. THEN ORDER THEM.

RECOMMENDATION RETRIEVAL AND RANKING



A useful recommendation must be relevant, available, and delivered in time.

Today: separate the shortlist from the final order, then measure each stage.

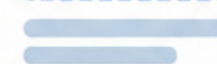
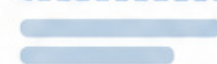
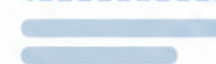


DEFINE THE RECOMMENDATION TASK

Illustrative task: show three related products on a product page.



You might also like



Define the user context, eligible catalog, relevance label, and response budget.



Similar items and personalized items are different tasks.

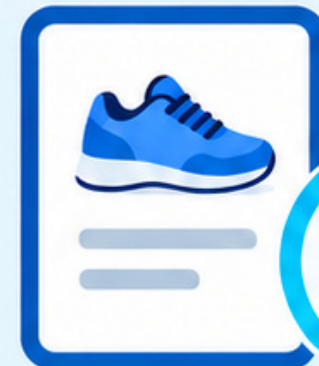


Similar items
(related to the current item)



Personalized items
(relevant to the user)

A related-item baseline can work without a user history.



WHY USE TWO STAGES?



RETRIEVAL finds a manageable set of plausible candidates.

RANKING spends more work ordering those candidates.



Wide catalog

Shortlist
(from retrieval)

Ranking

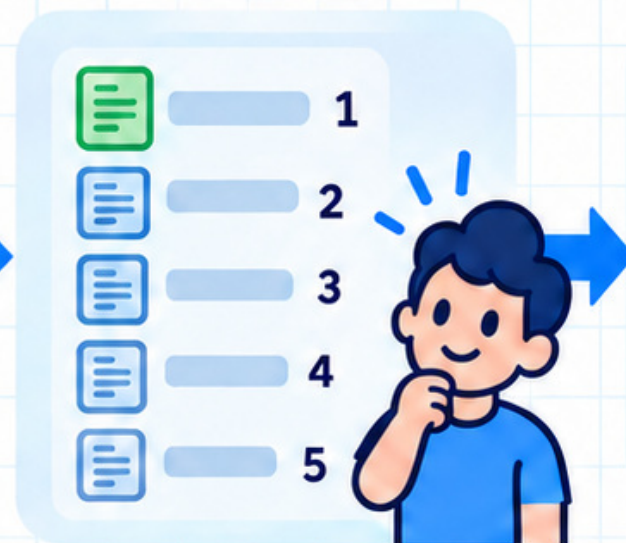
Final list



Many candidates



A manageable set of plausible candidates.



Top results



A relevant item missed by retrieval.

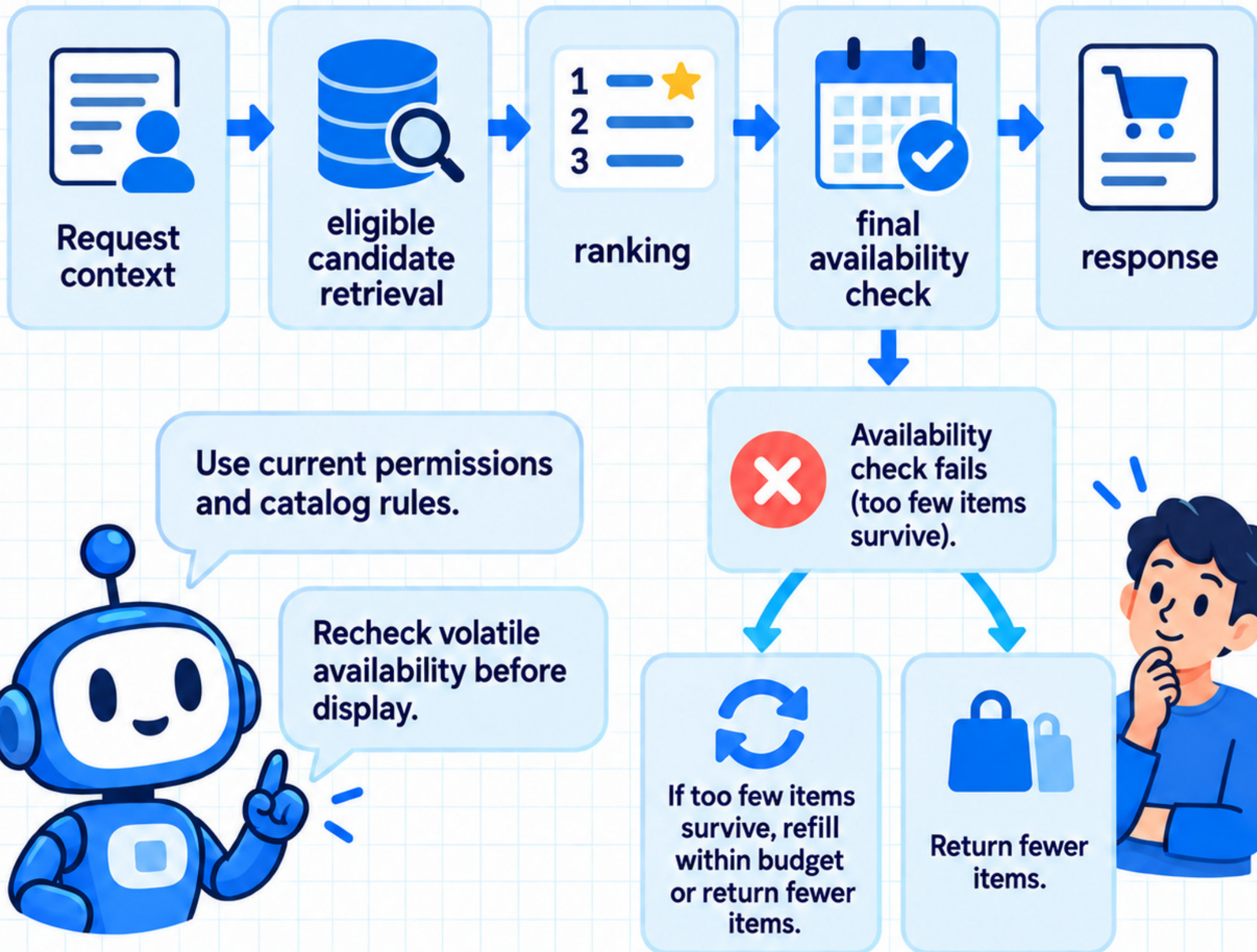


A ranker cannot recover a relevant item that retrieval never returned.



The candidate count trades coverage against downstream work.

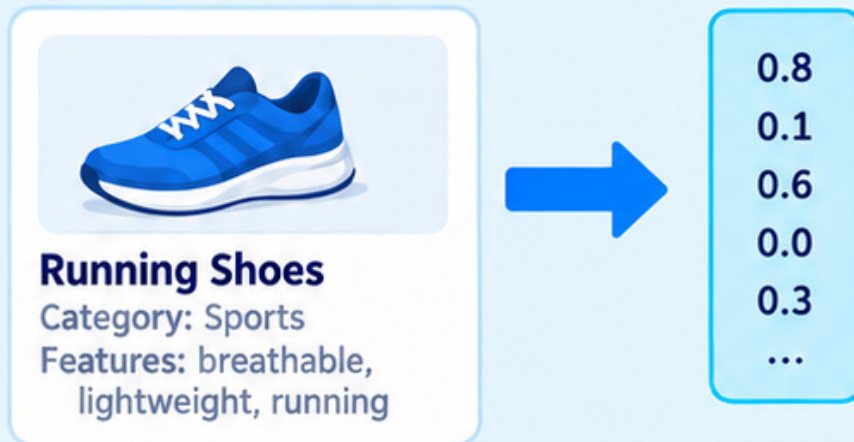
DRAW THE COMPLETE PATH



A SIMPLE RETRIEVAL BASELINE

Illustrative products, vectors, and configuration

- 1** Represent each product with declared text or category features.



- 2** Use nearest neighbors to find similar products.



- 3** Exclude the current product and duplicate IDs.



- 4** Version the representation and distance choice.

Representation: v1
Distance: cosine
Features: text + category



- 5** Similarity is a candidate signal, not proof that the user will buy.



WHAT THE TOOLS ACTUALLY DO

scikit-learn NearestNeighbors returns nearby rows under the configured distance.



Your application maps row indices to product IDs.

Illustrative mapping

Row Index	Product ID
0	P1001
1	P1043
2	P2098
3	P3175
...	...



PostgreSQL can filter catalog records using availability and access rules.



- ✓ In stock
- ✓ Meets access rules
- ✓ Visible in catalog

Application code combines candidates, applies budgets, and handles fallback.



Combine candidates

Apply budgets

Handle fallback

USE MORE THAN ONE CANDIDATE SOURCE



Illustrative sources:



Similar products from item features.



Popular eligible products in the category.



Personal-history candidates when appropriate data exists.



Merge by product ID and keep source provenance.

Deduplicated candidate product IDs with sources



ID: A1

Similar

Popular



ID: B2

Popular

History



ID: C3

Similar



ID: D4

Popular

History



ID: E5

History



A source quota can improve coverage, but must be evaluated.

RANKING HAS A DIFFERENT JOB

For each candidate, build features available at request time.



Estimate the declared relevance objective, then order the candidates.



A rule or simple model is a useful baseline.

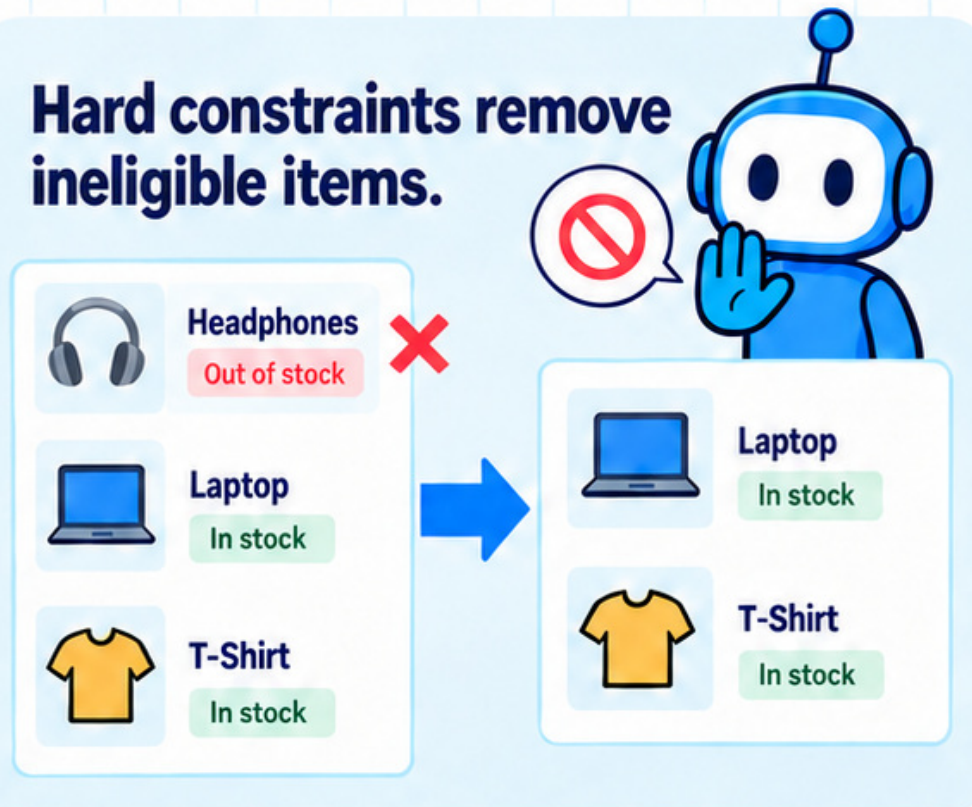


A high click score does not automatically mean high customer value.

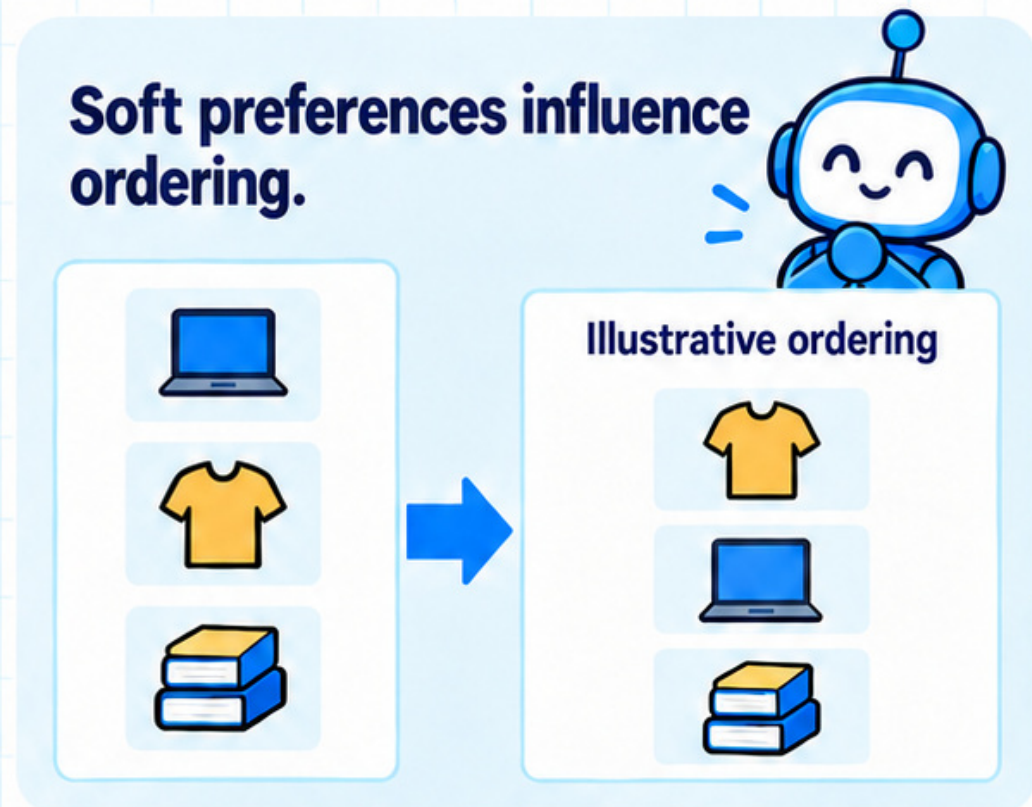


KEEP CONSTRAINTS EXPLICIT

Hard constraints remove ineligible items.



Soft preferences influence ordering.



Examples: unavailable stock is a hard exclusion; category variety can be a soft preference.

Apply the declared policy in code and evaluate the final list.

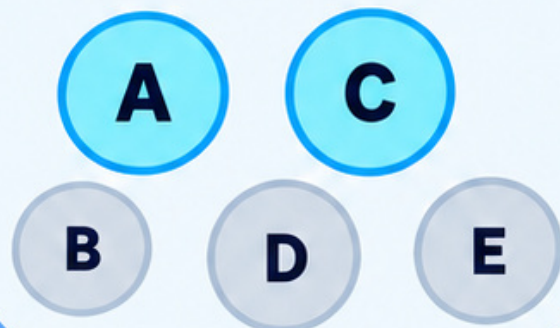


Do not hide business rules inside an unexplained score.

MEASURE CANDIDATE RECALL

Illustrative eligible relevant set: {A, C}.

All candidates
(in candidate universe)



Relevant items: {A, C}

Retrieved top 3: {A, C, D}.

Top 3 retrieved
(from candidate universe)



$\text{Recall@3} = \text{relevant items retrieved} / \text{all eligible relevant items}.$

Here: $2 / 2 = 1.00$.



Declare the candidate universe
and relevance labels before
comparing systems.



MEASURE THE FINAL ORDER

Illustrative binary relevance: A=1, C=1, D=0.

A

Relevance: 1

C

Relevance: 1

D

Relevance: 0



Ranking: D, A, C.

1**D**

Relevance: 0

**2****A**

Relevance: 1

**3****C**

Relevance: 1

$$\text{DCG@3} = 0 + 1/\log_2(3) + 1/\log_2(4).$$

Ideal order: A, C, D.

1**A**

Relevance: 1

**2****C**

Relevance: 1

**3****D**

Relevance: 0

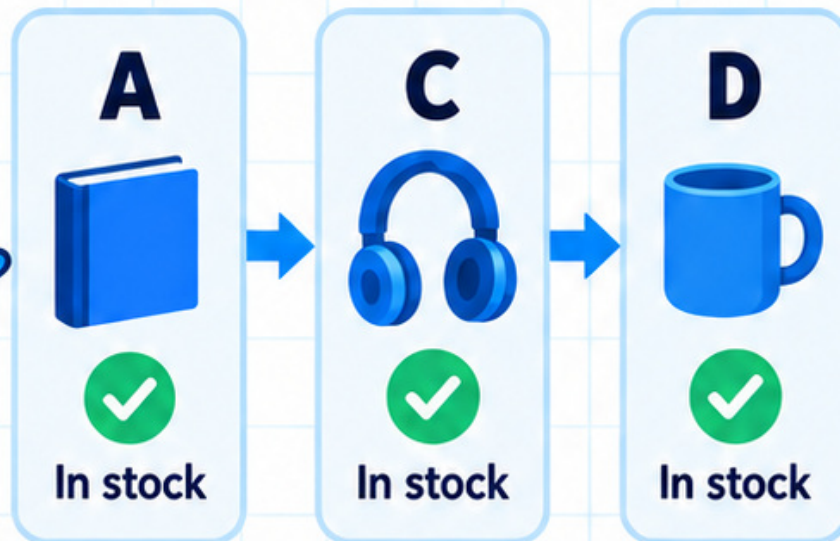
$$\text{NDCG@3} \approx 1.131 / 1.631 \approx 0.693.$$

Relevant items earn more credit near the top.



CHECK THE SERVED LIST TOO

Illustrative final list:
A, C, D.



FINAL CHECK RESULTS

Unavailable items
at the final check: 0.

Unavailable-item rate
 $= 0 / 3 = 0\%$.

0/3
= 0%



Also measure empty responses,
fallback rate, and latency.



A good ranking score cannot
excuse an ineligible recommendation.

COLD START NEEDS A ROUTE



FIT THE REQUEST BUDGET



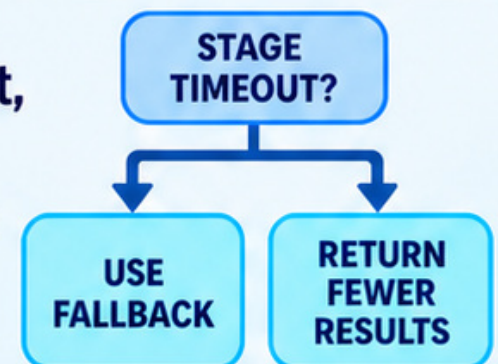
Allocate time to retrieval, feature reads, ranking, and final checks.

Measure the full request distribution, including fallbacks and failures.



Bound candidate count and refill attempts.

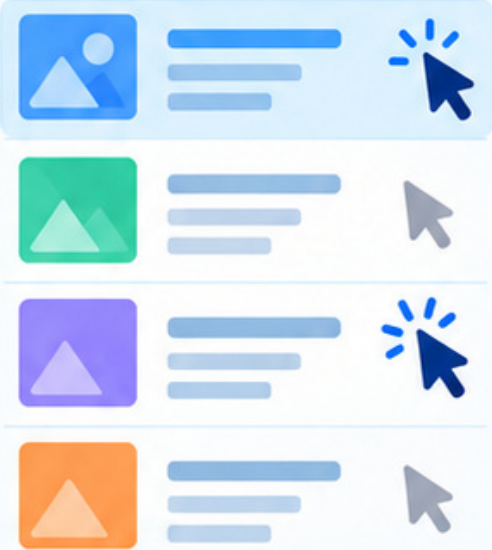
If a stage times out, use the declared fallback or return fewer results.



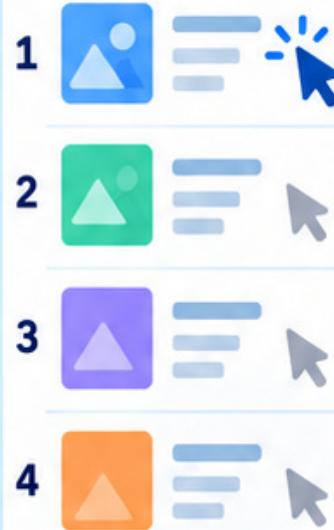
OFFLINE RELEVANCE HAS LIMITS

Clicks are observed only for displayed items.

Displayed to user



Position, exposure, and existing policy affect those clicks.



Position affects exposure.

Exposure affects clicks.

Existing policy affects which items are shown.

An unclicked item is not automatically irrelevant.

No click $\neq$ not relevant.



Not clicked (but could still be relevant)

Use a defined labeling protocol and realistic time splits.



Offline gains need careful online validation before business claims.



Offline gains



Careful online validation



Before business claims.

SAVE ENOUGH TO DIAGNOSE ERRORS



Log request context under an appropriate privacy policy.



Record candidate IDs and sources, feature and model versions, filter reasons, final order, and timing.



Join later feedback to the correct impression.

This separates a missing candidate from a bad ranking or availability failure.



FOLLOW ONE REQUEST



Retrieval

Candidate IDs and sources

Ranking

Feature and model versions

Filter

Filter reasons

Impression

Final order and timing

Later feedback

Join to the correct impression



YOUR TURN: LOCATE THE MISS



1 What is recall@3?

2 Which stage lost C?

3 Would a better ranker alone recover C?



EVALUATE EACH STAGE

Exercise: $\text{recall@3} = 1 / 2 = 0.50$.
Retrieval missed C; ranking cannot add it back.

RETRIEVAL (top 3)



MISSED

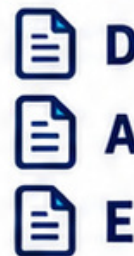


Not retrieved
→ cannot be
added back
by ranking.

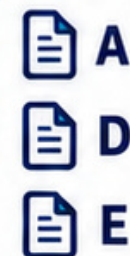
A ranker can move A upward,
while retrieval needs its own fix.

RANKING (reorder retrieved)

Before



After



RETRIEVAL FIX (add C)



Needs its
own fix.

Check relevance, availability,
coverage, and latency together.



Relevance



Availability



Coverage



Latency



Next: connect the
complete ML
serving design.

