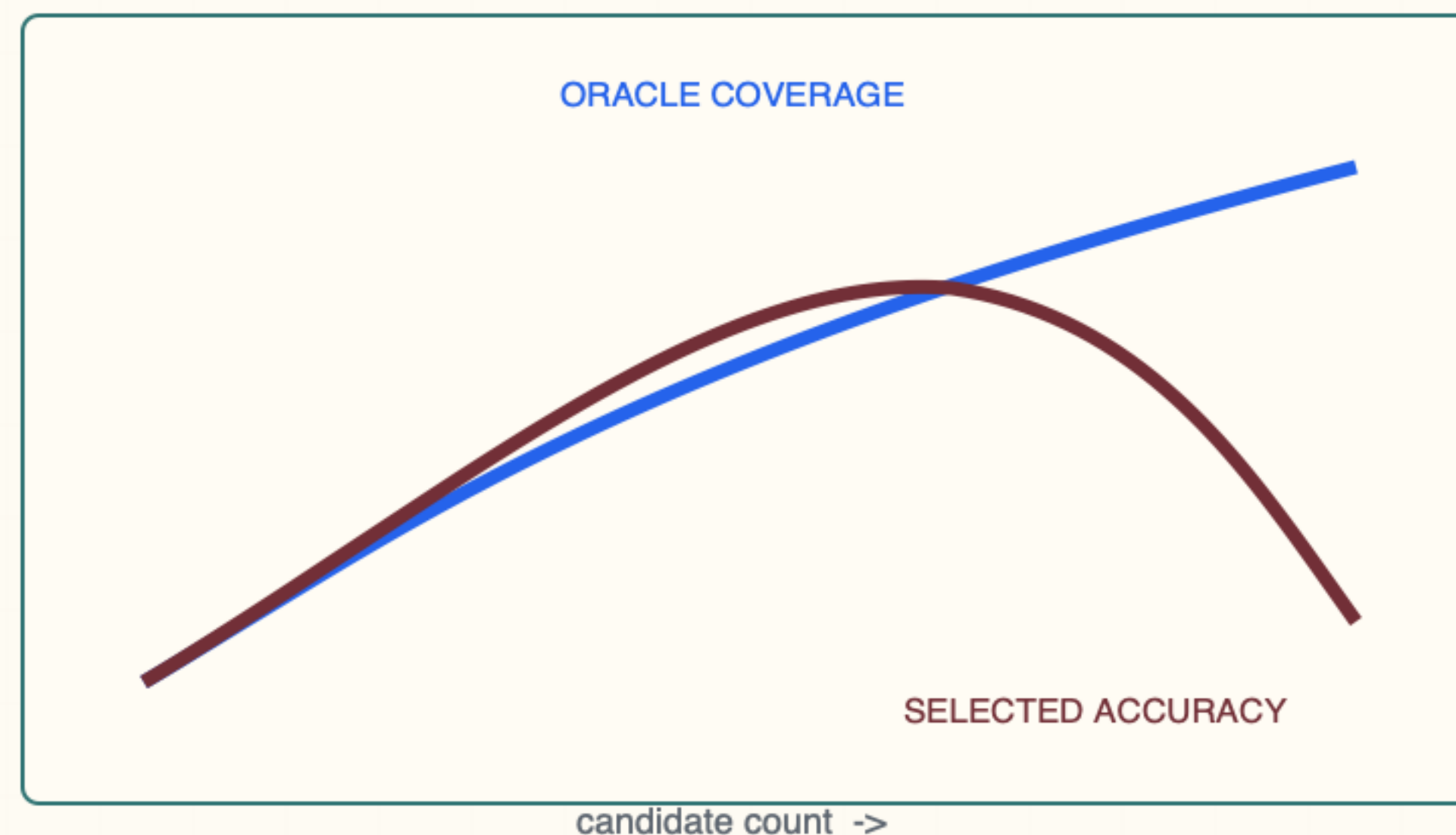


A Generator Is Only as Useful as Its Selector

Candidate coverage may rise while selected-answer accuracy stays flat or falls.

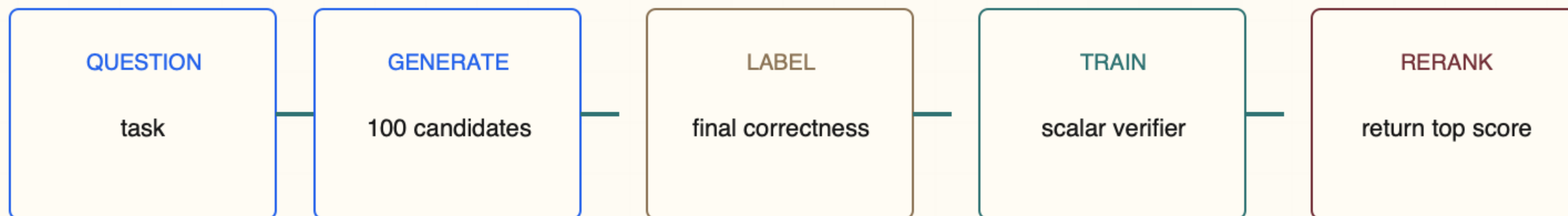
Generate more

Select one



The Basic Verifier Loop Learns to Rerank Completed Solutions

Generate, label, train a scalar verifier, then return the highest-scoring candidate.



CAVEAT

A learned correctness score is evidence for ranking, not proof that the selected answer is true.

More Candidates Create More False-Positive Opportunities

The maximum noisy score becomes less trustworthy unless verifier precision scales too.

AS N GROWS

chance of a correct candidate rises

chance of an extreme verifier false positive rises too

NOISY SCORE COMPETITION



green exists $\neq$ red wins

CAVEAT

The roughly 400-candidate peak shown in the lecture is setup-specific, not a universal optimum.

Outcome Scores Judge the Finish; Process Scores Judge the Route

PRMs add step-level credit assignment to a final-answer signal.

OUTCOME REWARD MODEL



final: 0.78

PROCESS REWARD MODEL



0.9

0.3

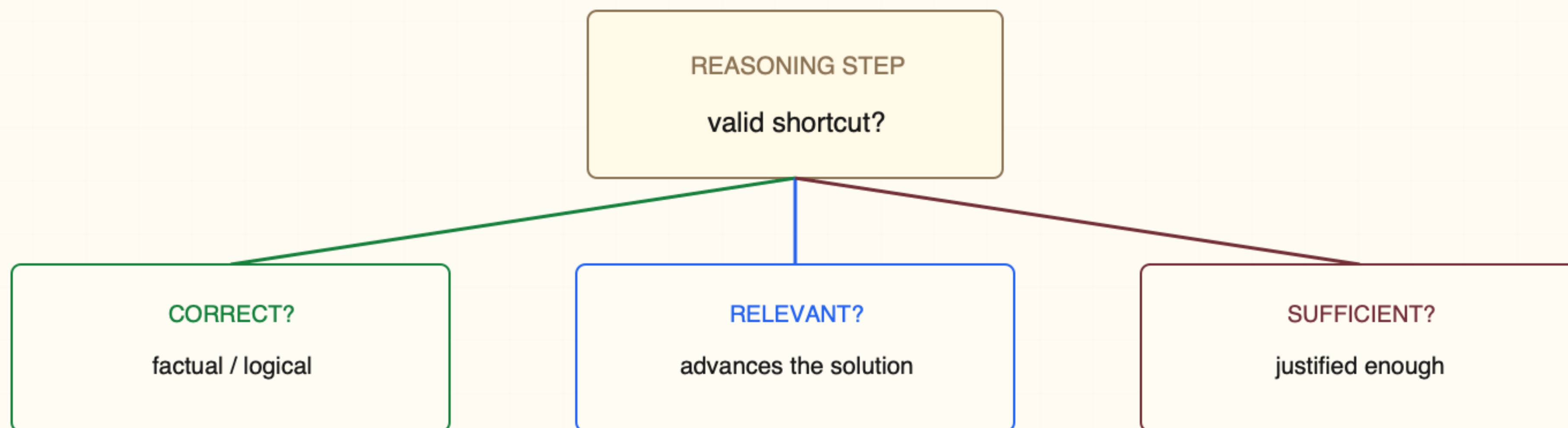
0.8

0.7

Terminal credit <-> step-level credit assignment

Process Labels Are Richer, but They Encode a Rubric

Correctness, relevance, and sufficient justification depend on an annotation policy.

**CAVEAT**

A process label encodes an annotation policy and may reject valid alternative reasoning styles.

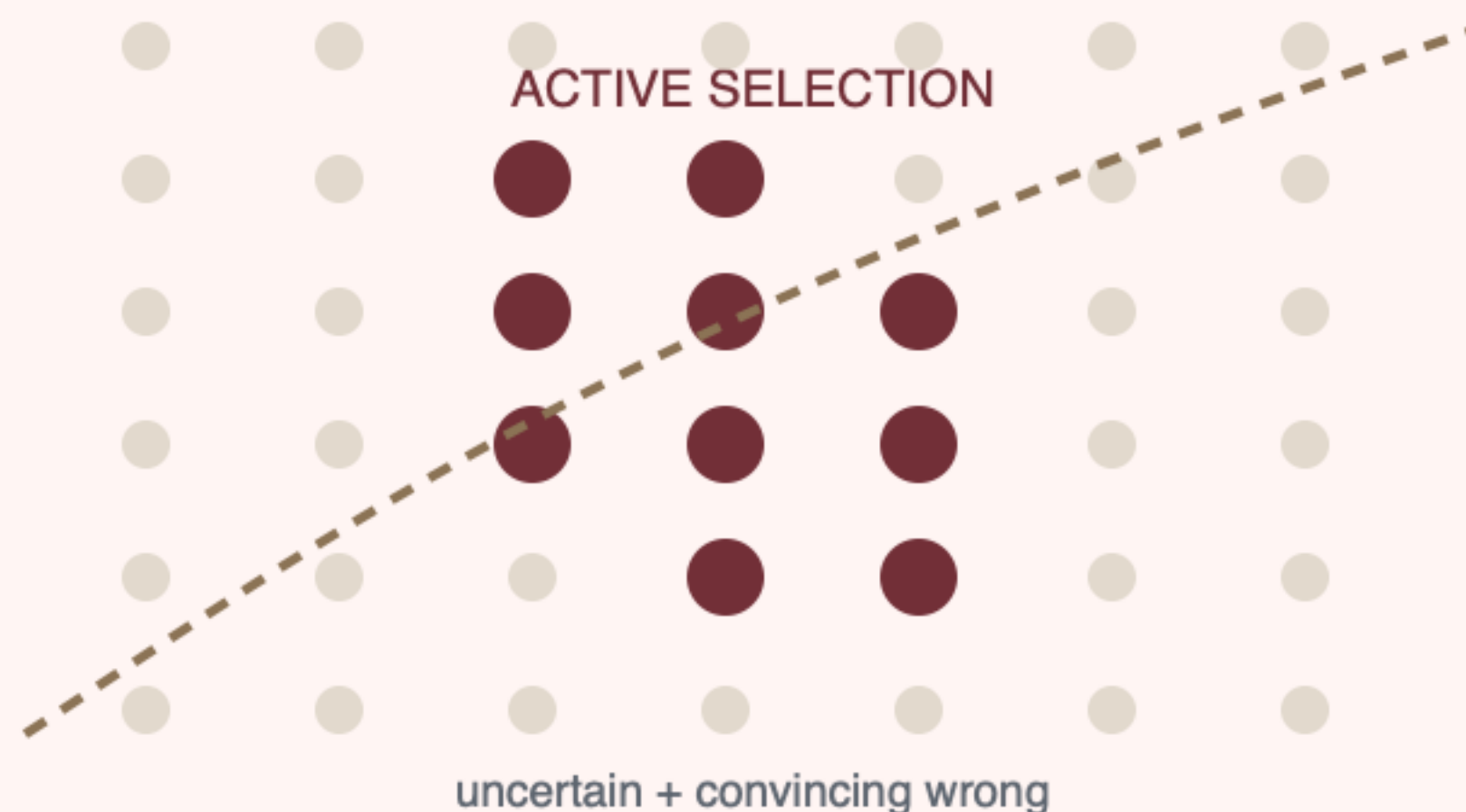
PRM800K Targets Informative Mistakes Instead of Labeling Uniformly

Active selection concentrates annotation near uncertain and convincing wrong solutions.

UNIFORM LABELING



ACTIVE SELECTION

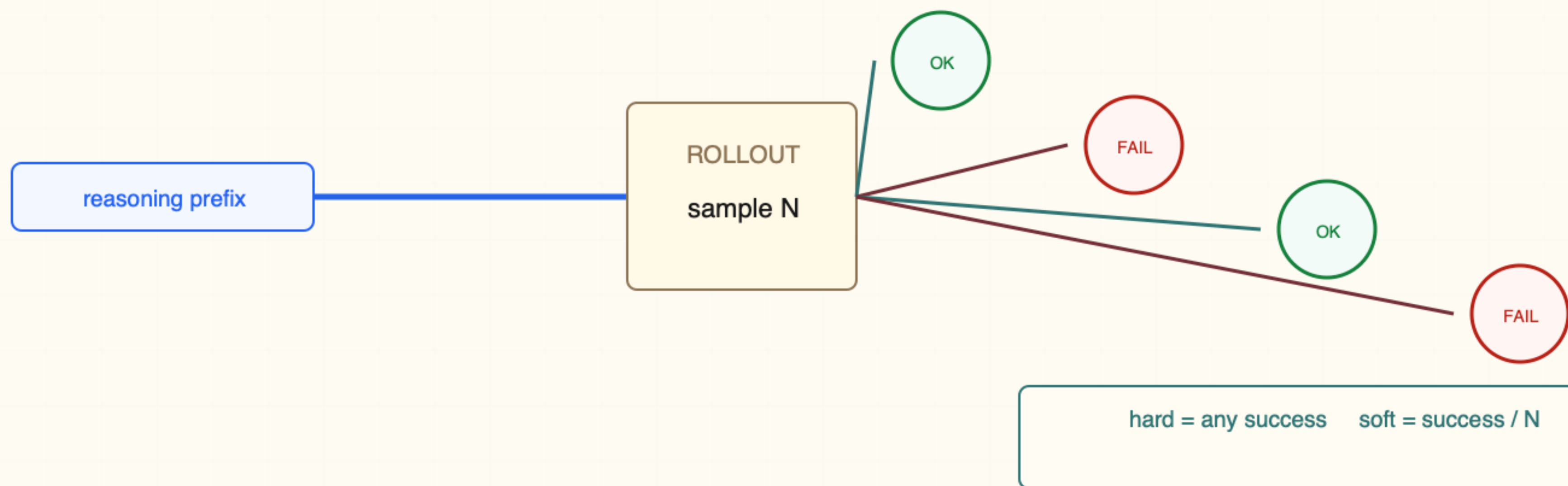


CAVEAT

Active selection changes the training distribution and complicates direct label-budget comparisons.

Automatic Process Labels Trade Human Cost for Rollout Compute

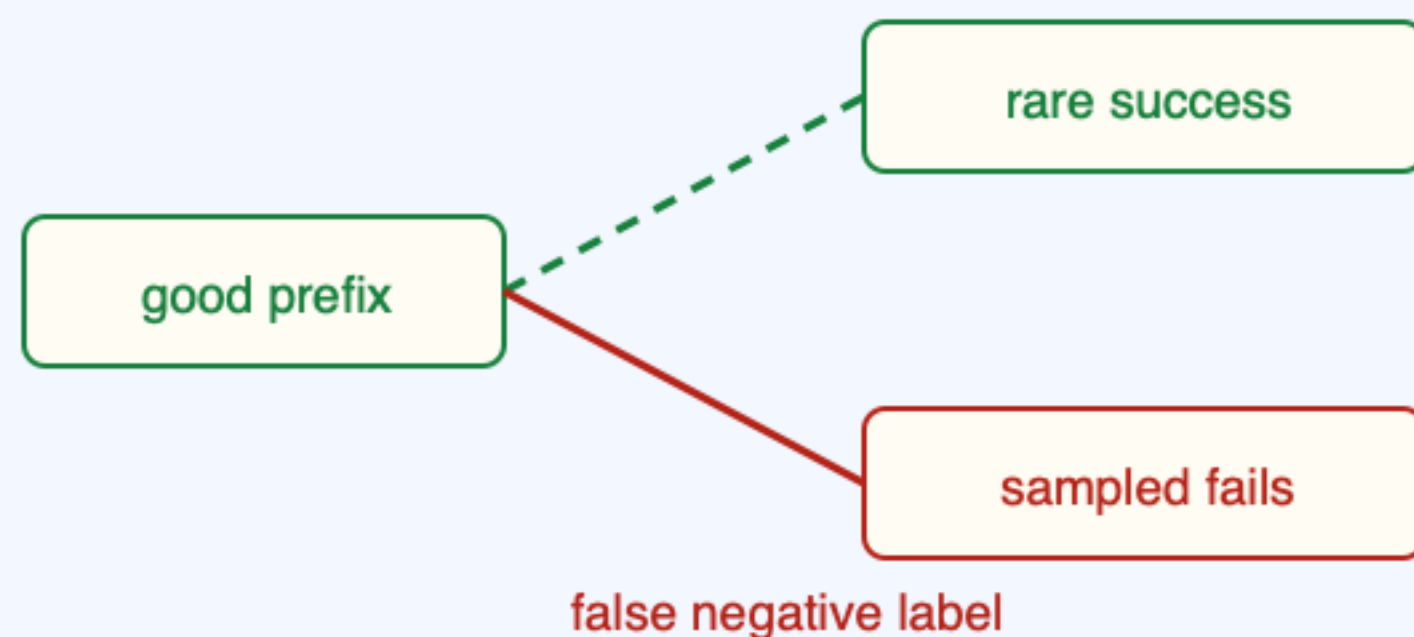
A prefix is judged by sampled continuations that reach the correct final answer.



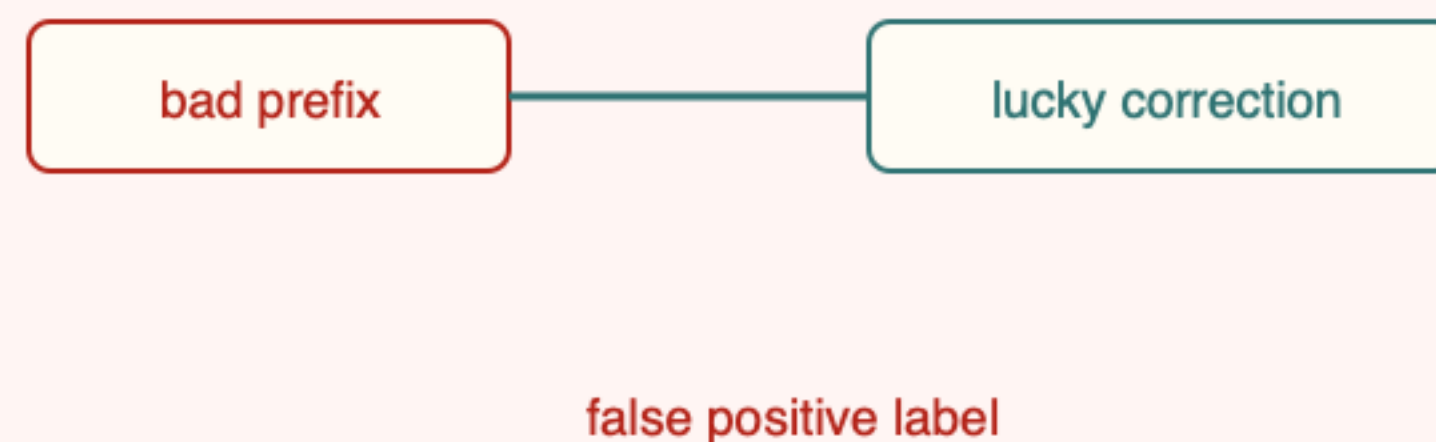
Rollout Labels Inherit the Search Policy's Blind Spots

Good rare paths can look bad; wrong steps can look good after a lucky repair.

RARE VALID PATH MISSED



WRONG STEP LUCKILY REPAIRED

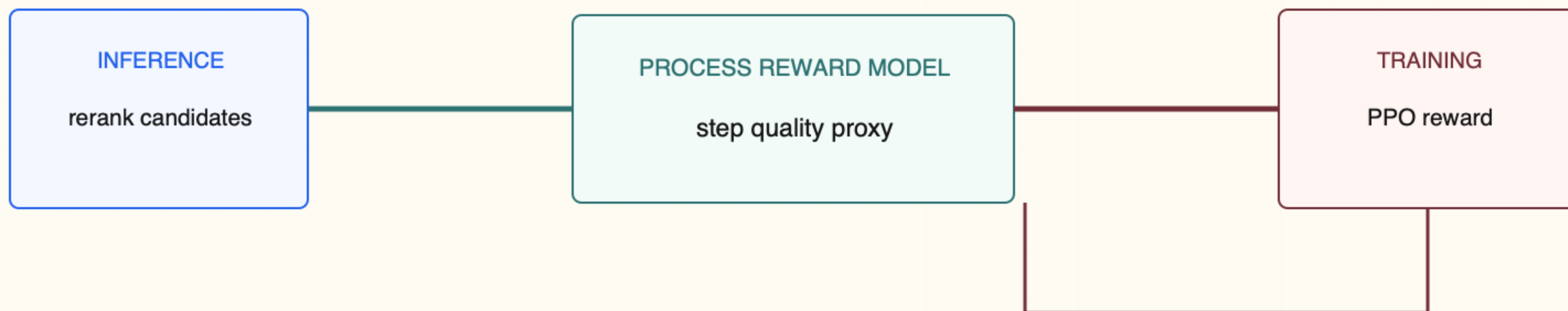


CAVEAT

Rollout labels mix prefix quality with the capability and luck of the rollout policy.

Verification Can Select Outputs and Train the Generator

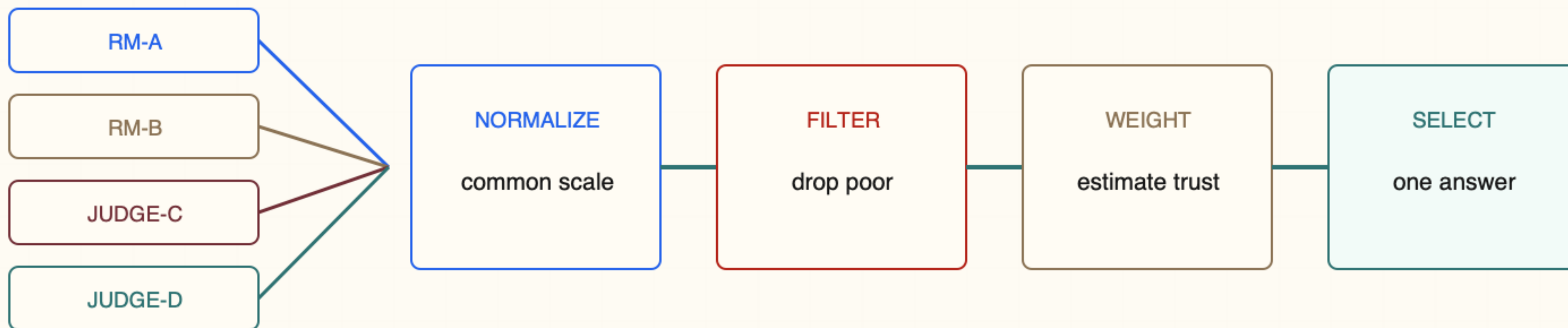
The same PRM can guide inference-time reranking and PPO-style policy updates.

**CAVEAT**

Optimizing an imperfect reward model can raise the proxy score without improving correctness.

Diversity Across Verifiers Helps Only After Weighting and Filtering

Weaver normalizes heterogeneous scores, removes poor verifiers, and combines weighted evidence.



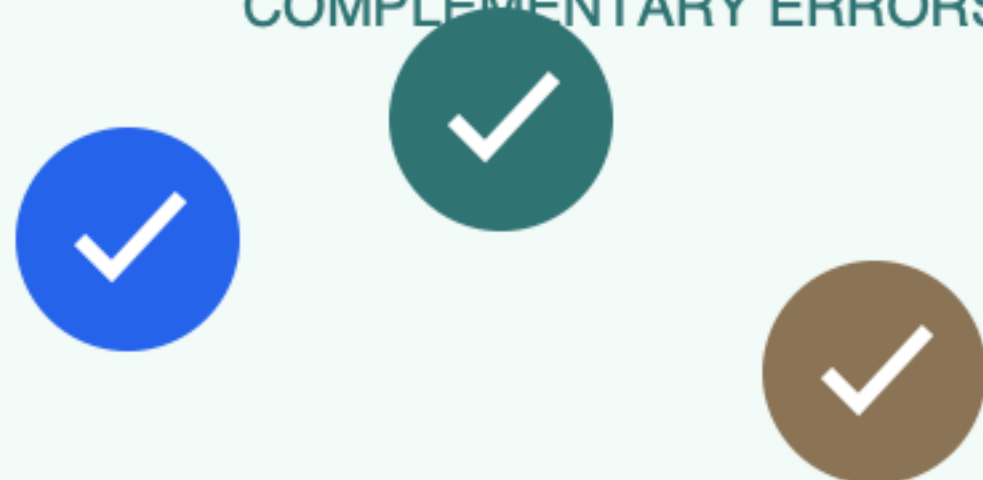
CAVEAT

Adding more verifiers naively is not monotonically useful;
sources differ in quality and bias.

Shared Blind Spots Can Defeat an Entire Verifier Ensemble

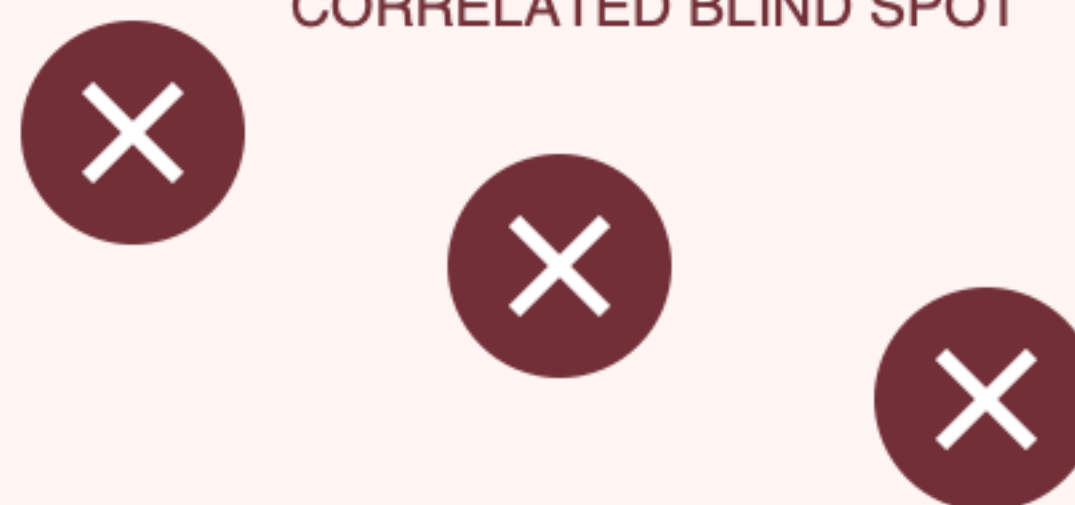
Complementary errors add evidence; correlated errors repeat the same mistake.

COMPLEMENTARY ERRORS



different evidence adds value

CORRELATED BLIND SPOT



same mistake repeated

CAVEAT

Weighting cannot create independent evidence when every verifier shares the same failure.

Design Verification as a Budgeted End-to-End System

Evaluate selection accuracy, calibration, transfer, cost, and coverage together.



CANDIDATES

coverage



GENERATOR

capability



VERIFIERS

discrimination



DISTILL

serving cost

SELECTED-ANSWER ACCURACY + CALIBRATION + COST

NEXT: PART 4 - LEARNING FROM FEEDBACK WITH TOOLS AND CODE