

More Inference Compute Creates Options, Not Guarantees

Stanford CS329A Part 2 - generation,
allocation, and verification under a fixed
model

**FIXED
MODEL**
weights locked

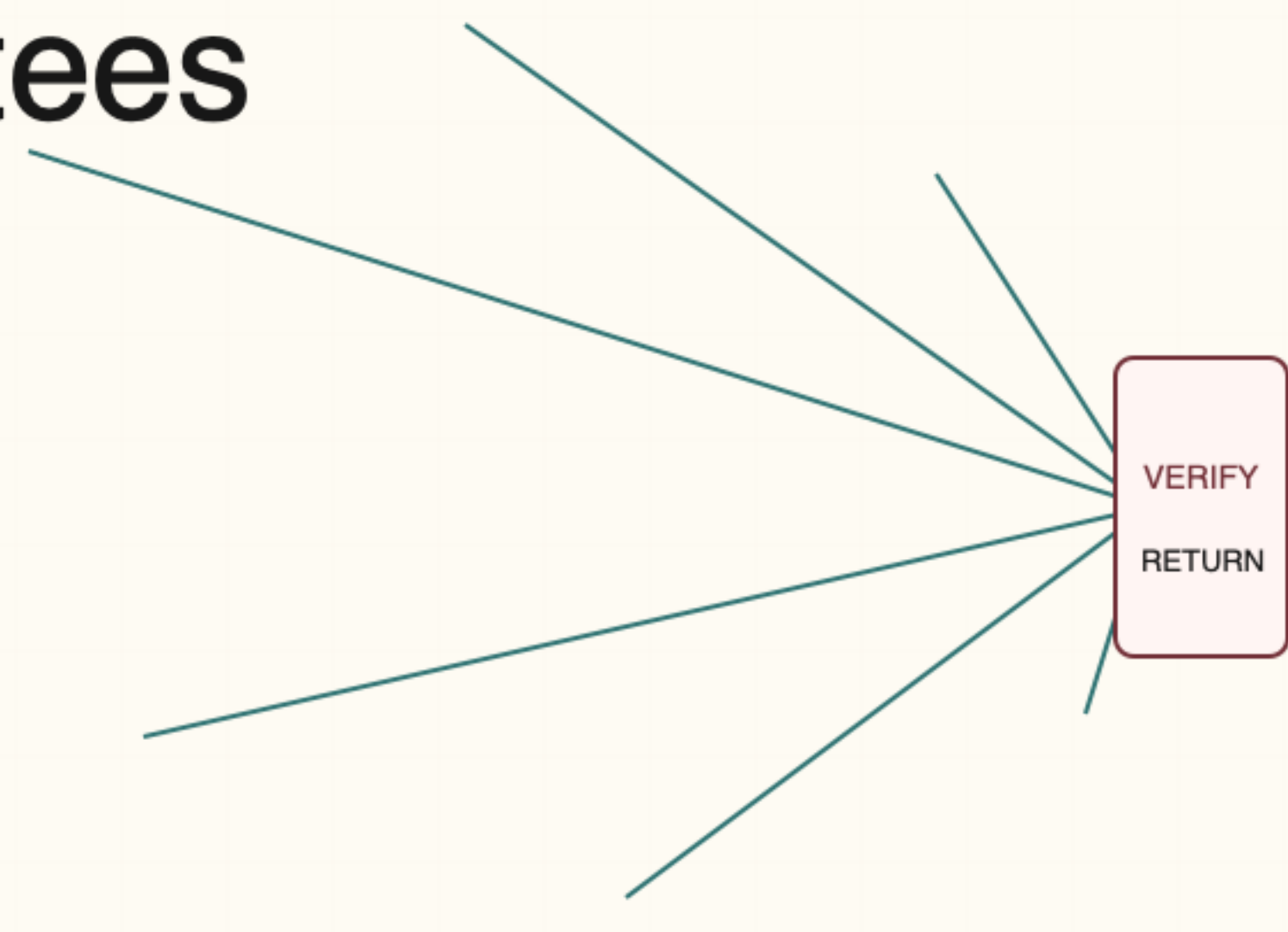
Generate

Allocate

Verify

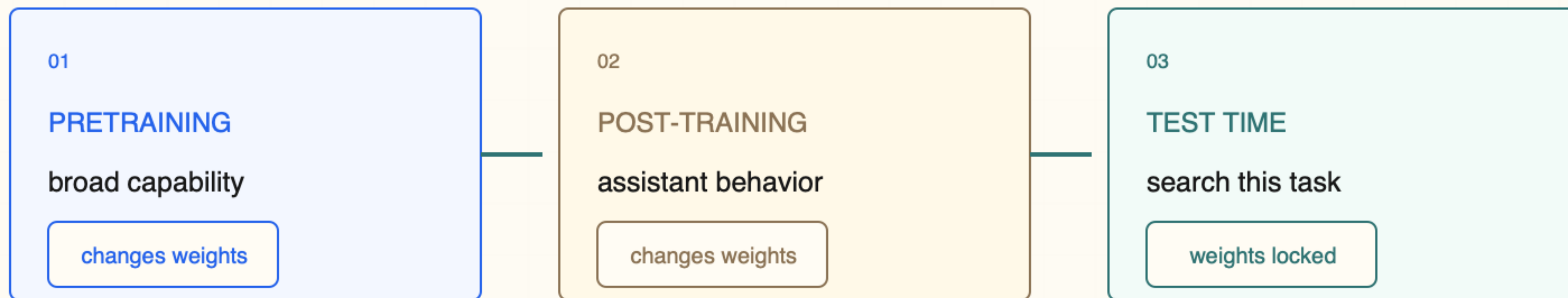
Stop

VERIFY
RETURN



Test-Time Scaling Searches Without Changing the Weights

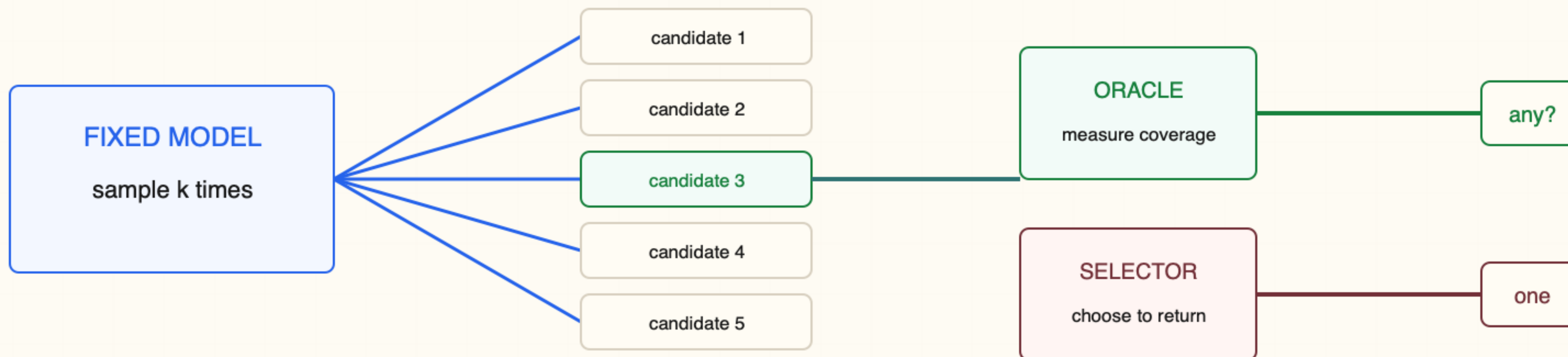
Training changes the model; inference spends a recurring per-request budget.



TRAIN ONCE -> AMORTIZE | INFER PER REQUEST -> PAY AGAIN

Repeated Sampling Buys More Chances

One fixed model produces many candidates; a verifier determines whether the gain is usable.

**CAVEAT**

Oracle coverage can rise even when a deployable selector cannot find the correct candidate.

A Correct Candidate Can Exist and Still Never Reach the User

Coverage belongs to the candidate set; reliability belongs to the whole pipeline.

COVERAGE / PASS@K

Does any candidate contain the correct answer?

candidate-set metric

RETURNED RELIABILITY



green exists != red selected

CAVEAT

A benchmark-to-product claim must name both the candidate metric and the selection mechanism.

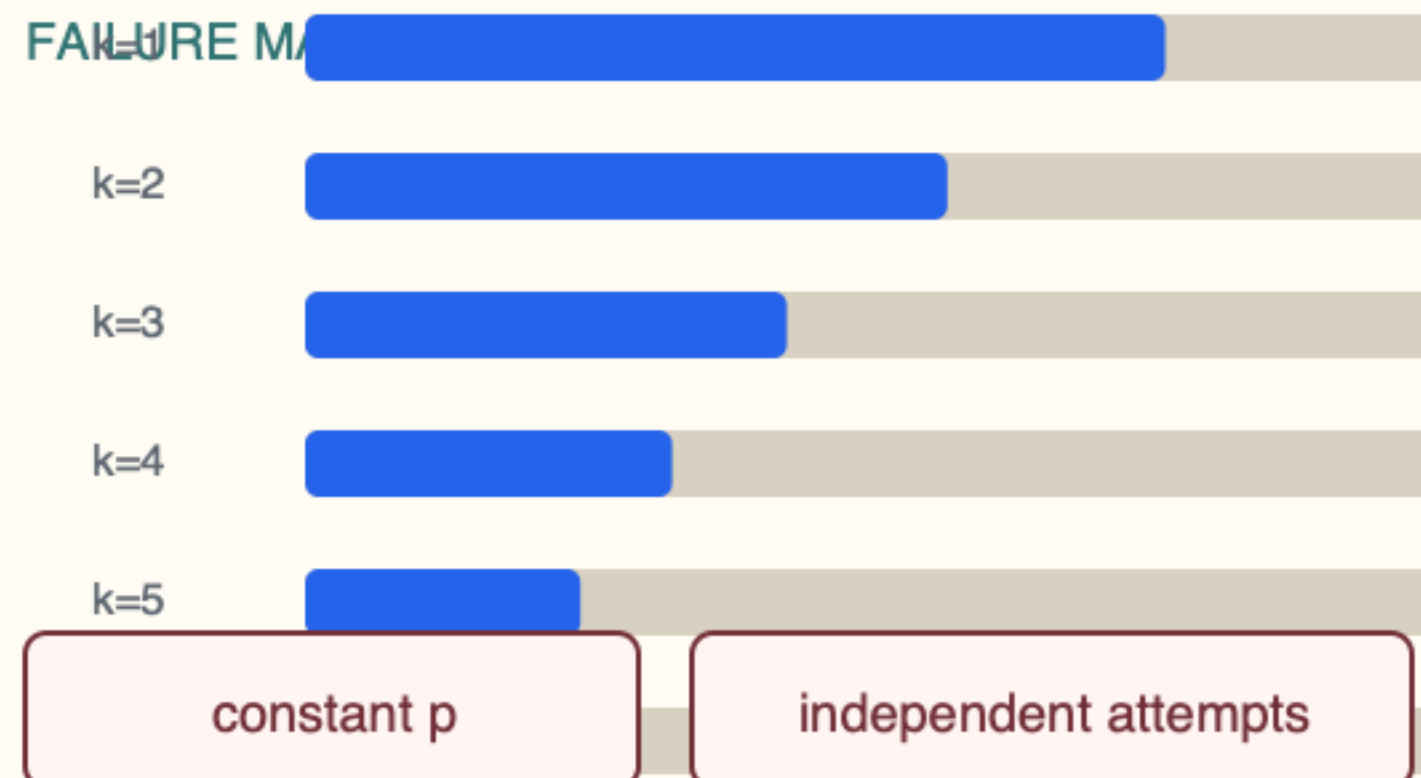
One Problem Has an Exponential Failure Curve

The familiar pass@k formula is an idealized model with visible assumptions.

pass@k

$$1 - (1 - p)^k$$

failure after k = $(1 - p)^k$

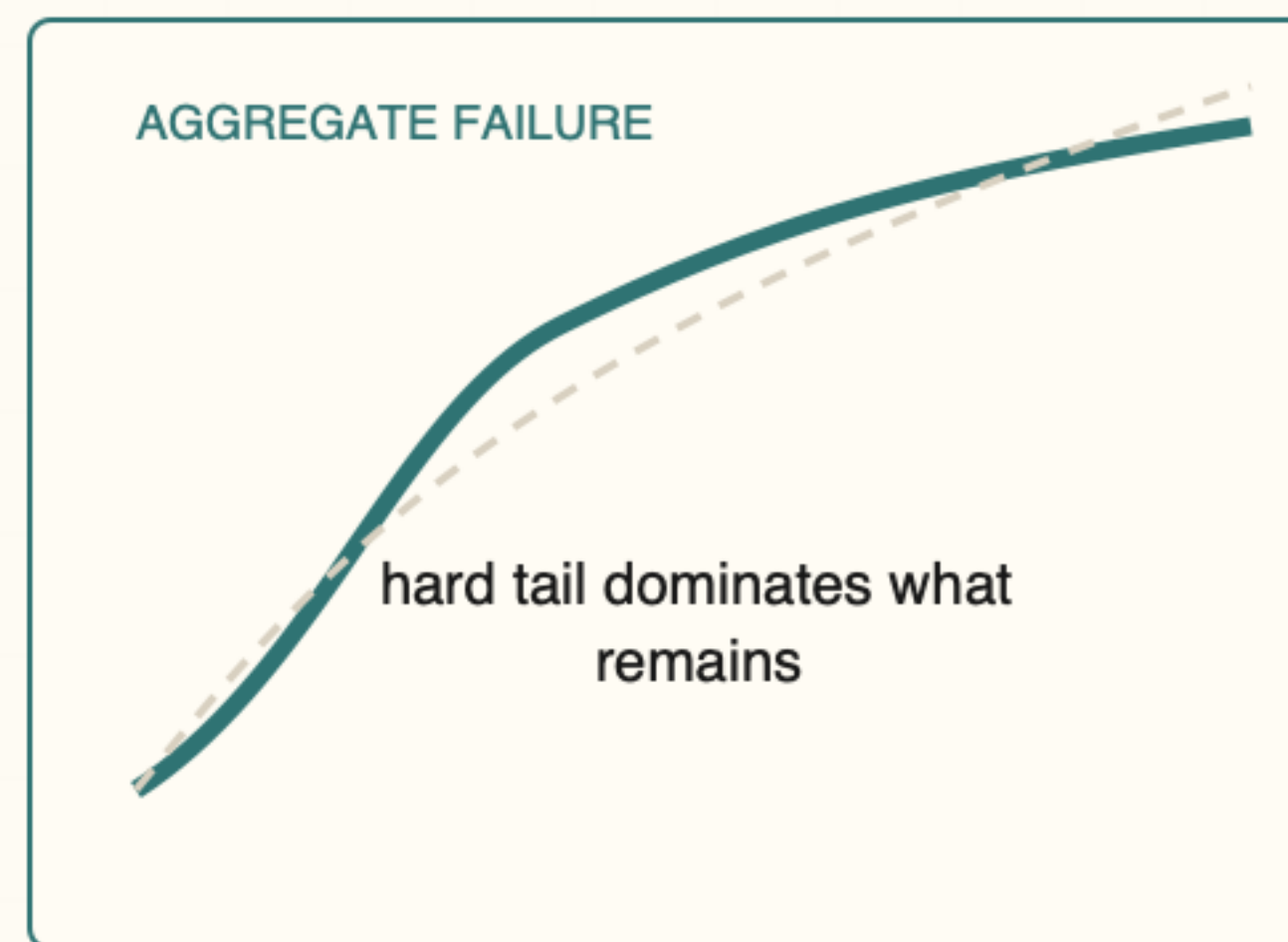
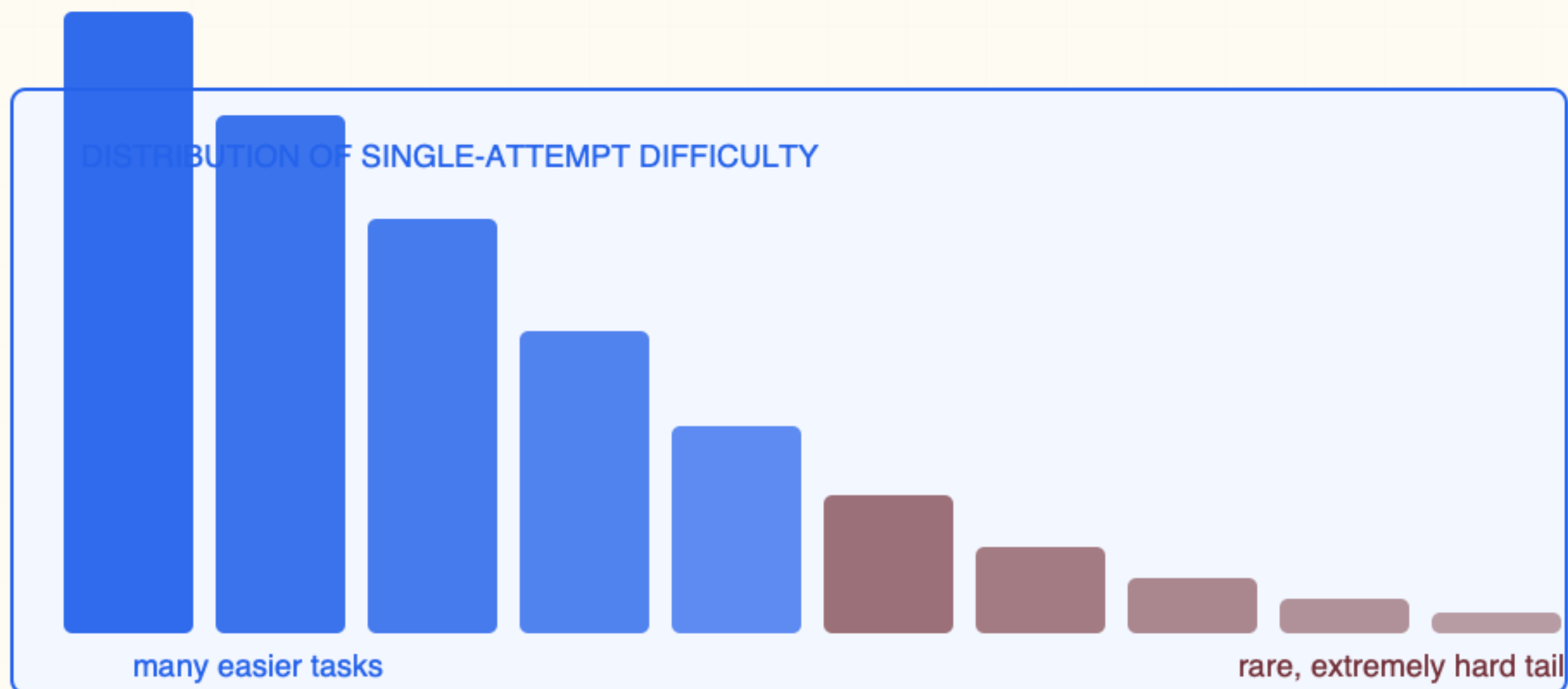


CAVEAT

Samples from one model can be correlated, so effective diversity may be much lower than k.

A Few Extremely Hard Problems Bend the Benchmark Curve

Per-problem exponential scaling and aggregate power laws can coexist.



CAVEAT

The aggregate curve can look power-law scaled even though each fixed problem decays exponentially.

Generation Scales Only as Far as Verification Can Follow

More candidates help when evidence can distinguish success from polished error.

CODE

execute tests



verification leverage

FORMAL PROOFS

proof checker



verification leverage

MATH

rules / exact answer



verification leverage

WRITING + SCIENCE

expert judgment

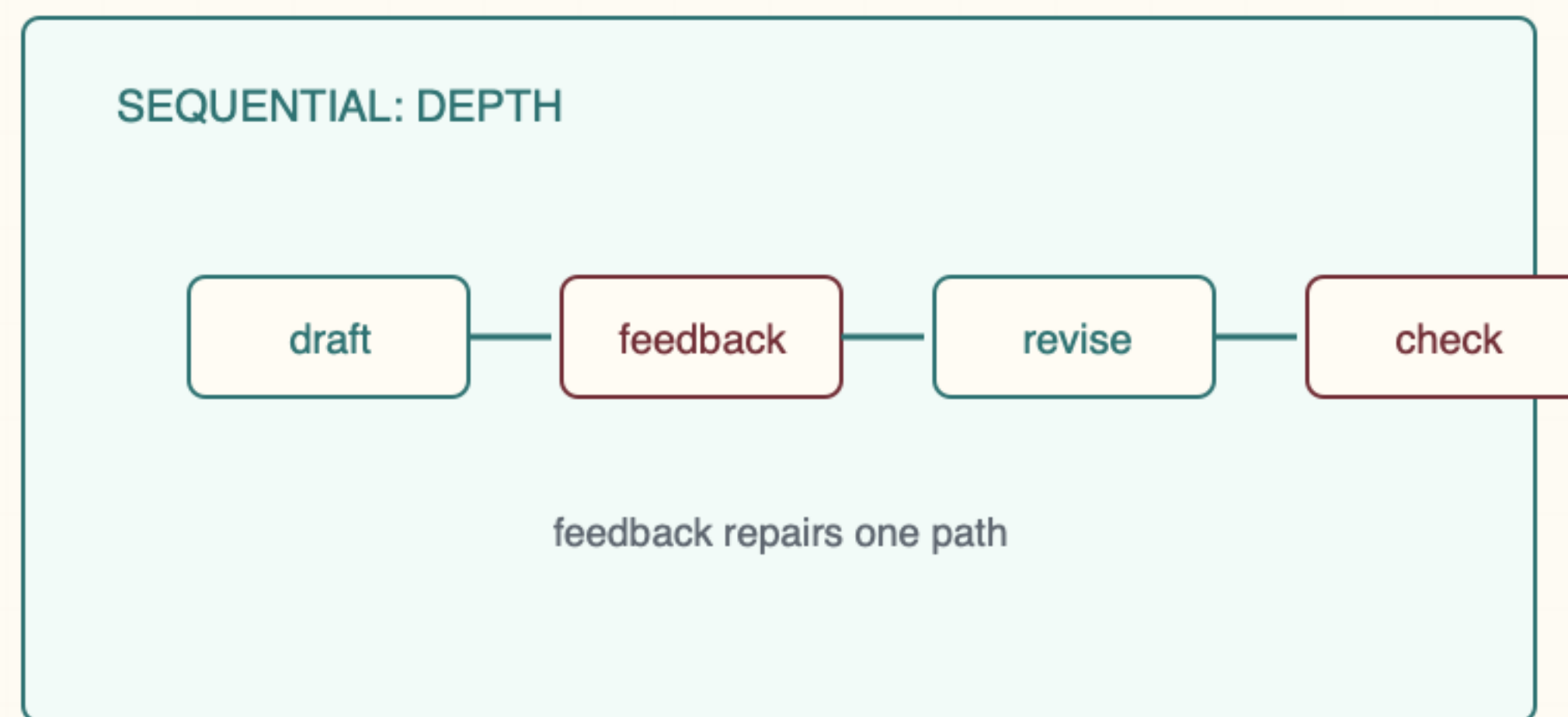
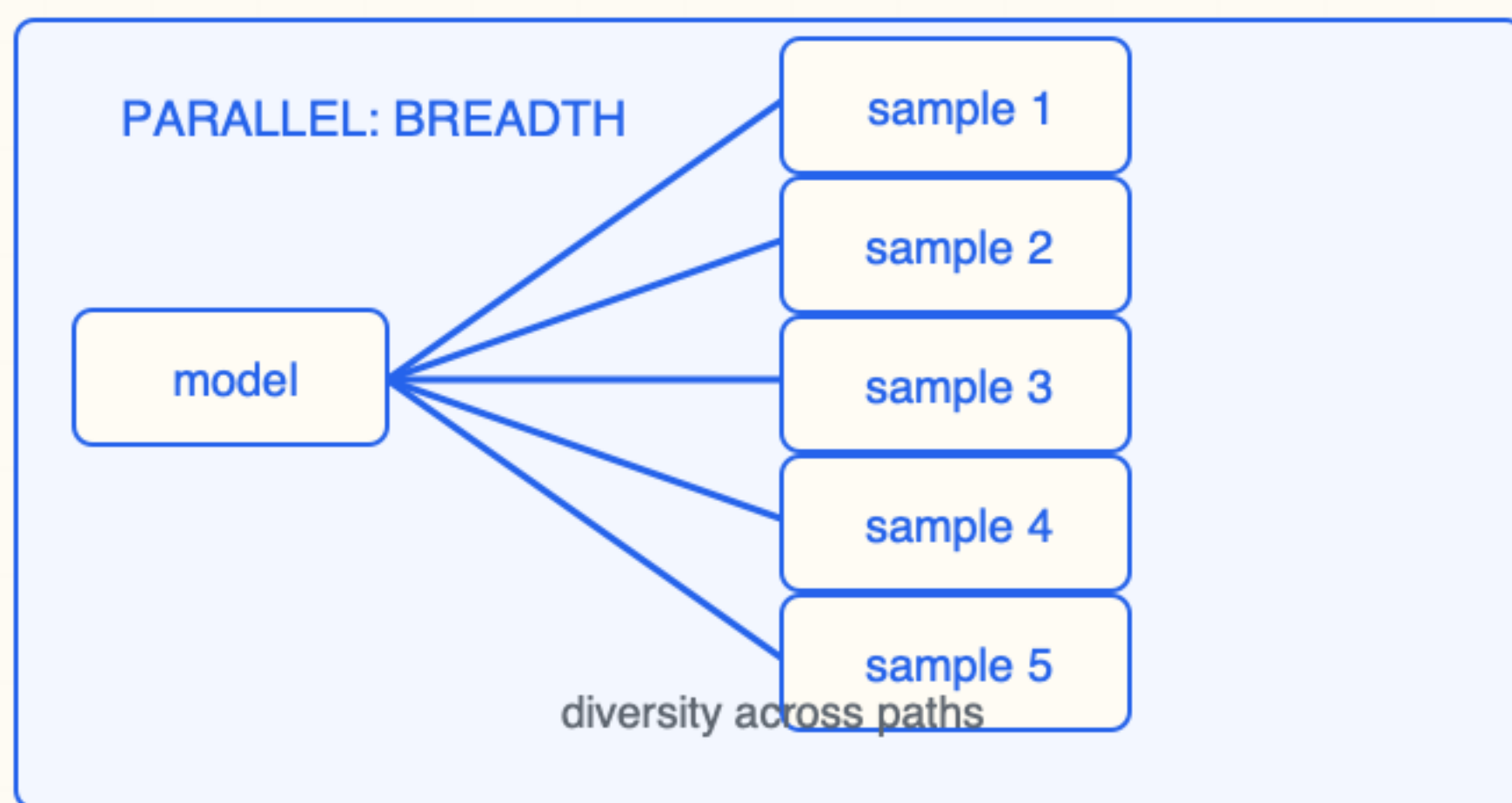


verification leverage

more judgment, weaker automation, higher feedback cost

The Same Budget Can Explore Breadth or Refine Depth

Parallel samples diversify; sequential revisions use feedback to repair a path.



SAME TOTAL BUDGET - DIFFERENT ALLOCATION

Outcome Scores Choose Finishes; Process Scores Shape the Route

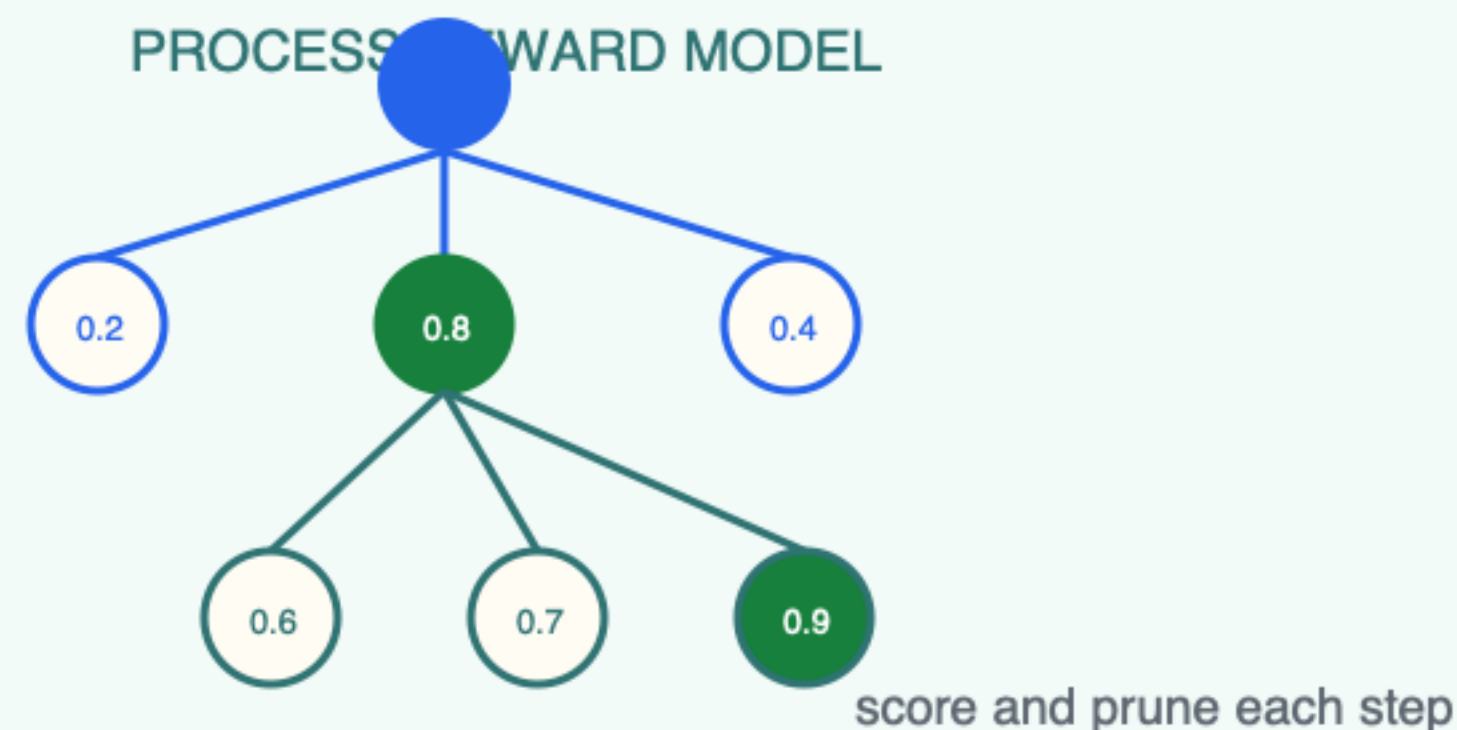
Reward models guide search, but neither becomes ground truth by definition.

OUTCOME REWARD MODEL



scores the finish

PROCESS REWARD MODEL

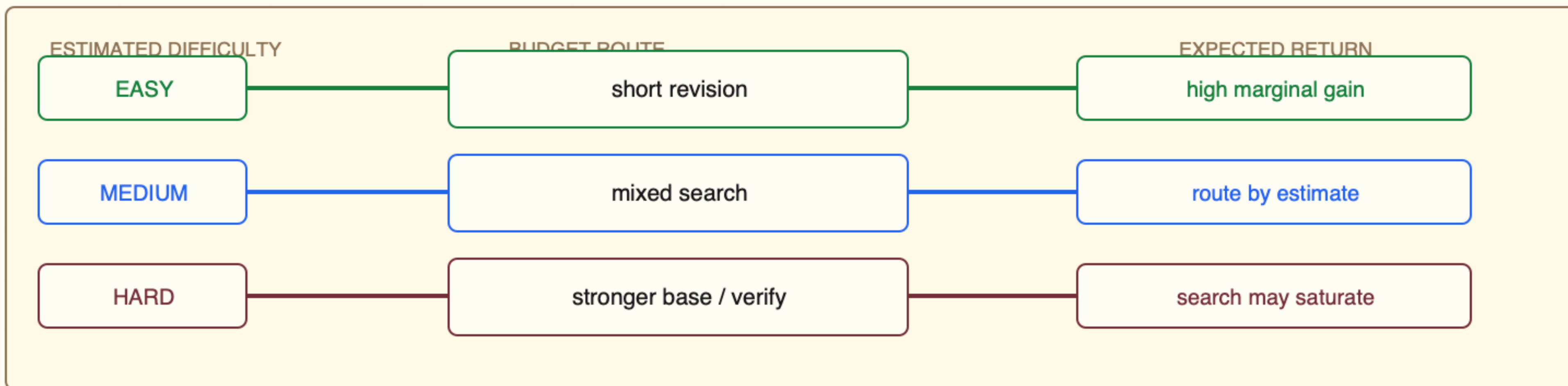


CAVEAT

Learned reward models can be confidently wrong or fail outside their training distribution.

Compute Allocation Should Change With Estimated Difficulty

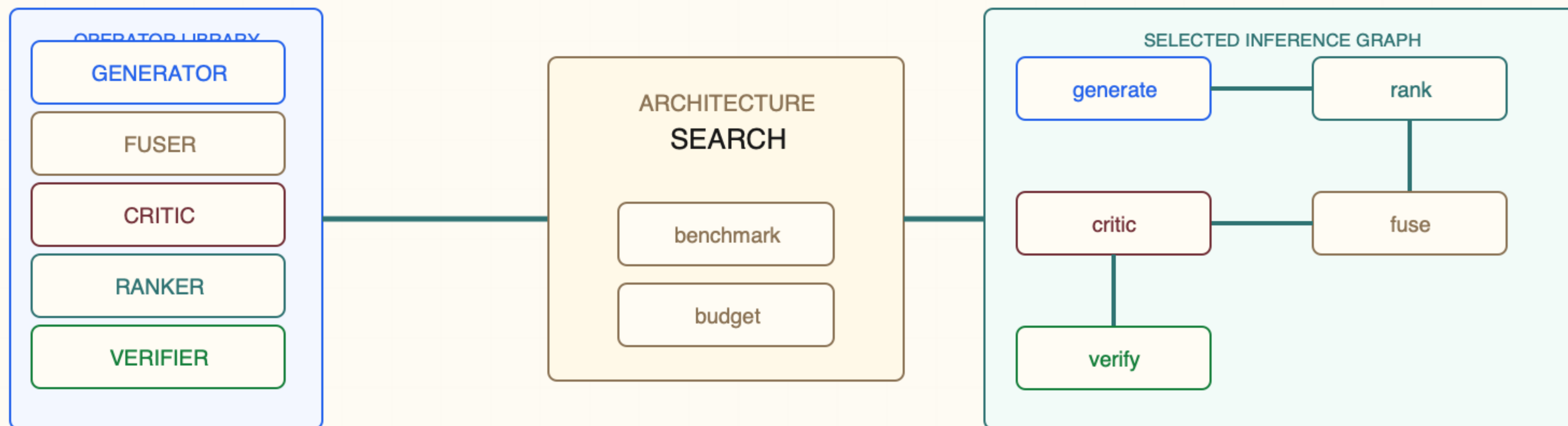
A uniform best-of-N policy can waste budget on tasks with different marginal returns.

**CAVEAT**

Compute-optimal means optimal for the chosen model, verifier, methods, difficulty estimate, and budget.

Archon Searches Over Inference Architectures, Not Just Answers

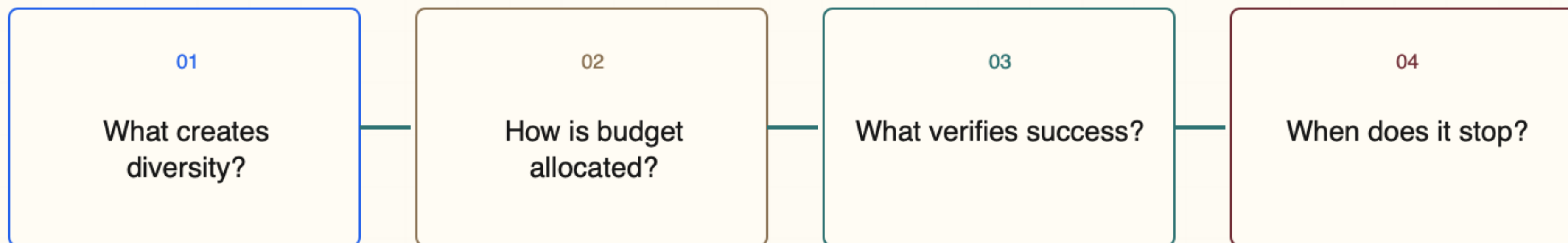
Models and inference operators become components in a budgeted architecture search.

**CAVEAT**

Results are benchmark-, version-, operator-set-, and budget-specific.

Before Buying More Inference, Find the Bottleneck

Use four questions to audit any test-time scaling proposal.



NEXT: PART 3 - ROBUST VERIFICATION

The bottleneck moves when generation, selection, cost, or latency saturates.