



## 1. Problem & Task Decomposition

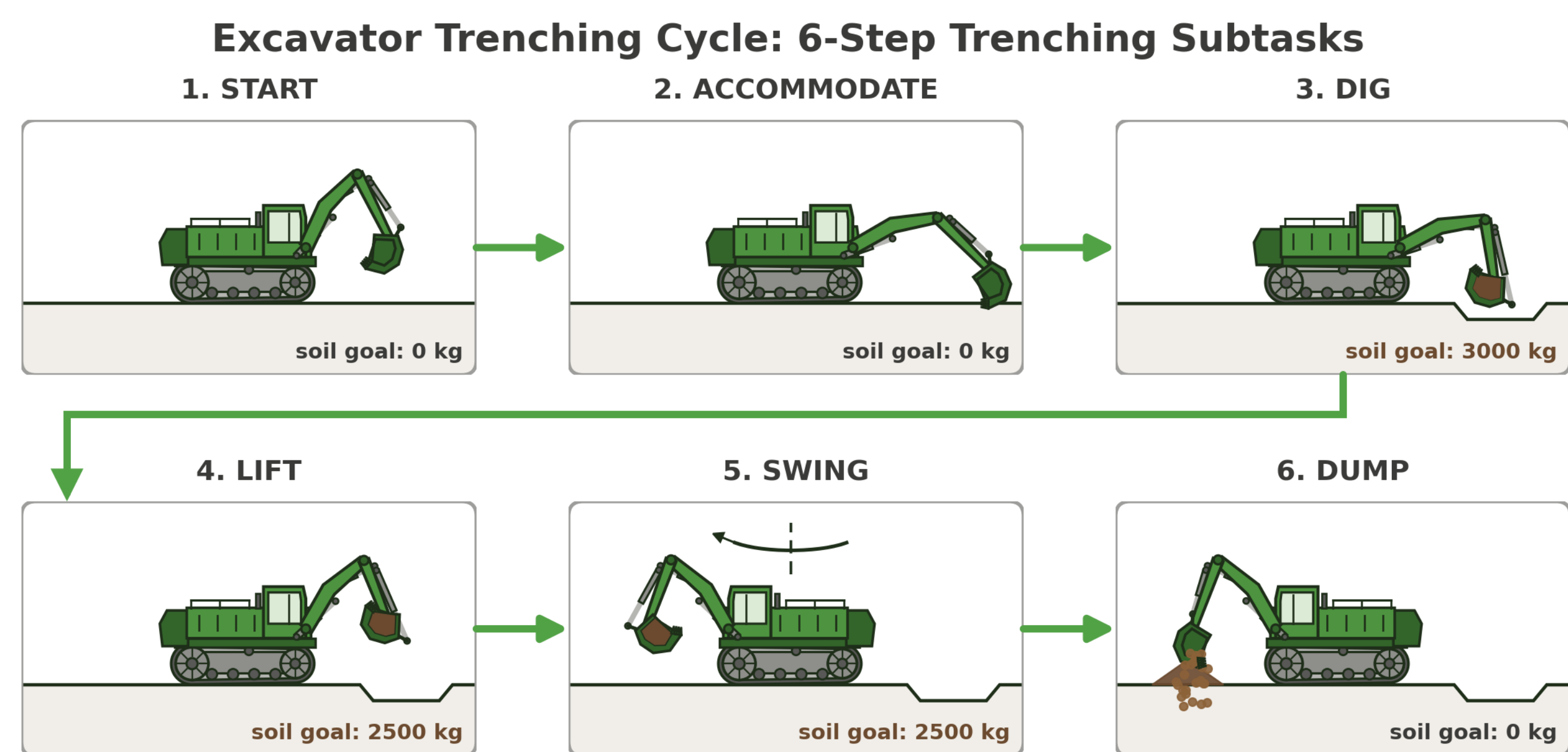
**We examine excavator trenching for three reasons:** [6] it is a representative earthmoving task; soil responds far less predictably than the arm's own geometry; and the horizon is long — what the bucket takes at the dig phase sets the load for every stage after it, and with it the number of cycles needed to finish the trench.

**Neither approach survives on its own.** Model-free RL sees reward long after the actions that earned it — *reward sparsity*, credit assignment failure, divergence. Classical control is untroubled by the horizon, but cannot anticipate the soil.

|             | Classical Control (IK)                                         | Model-Free RL                                                                               |
|-------------|----------------------------------------------------------------|---------------------------------------------------------------------------------------------|
| <b>Pros</b> | Predictable, smooth trajectory tracking in free space.         | Learns complex, non-linear dynamics without manual gain tuning or explicit system modeling. |
| <b>Cons</b> | Approximate unless tuned per machine, attachment and scenario. | Highly unstable when unguided; prone to cold-start exploration failures.                    |

**Core research question.** Can Classical Control Guidance (via Curriculum Learning & Residual RL) accelerate Model-Free RL training and yield superior task execution on earthmoving tasks?

**6-Stage Decomposition.** Excavation cycles natively follow a sequential kinematic sequence (Start, Accommodate, Dig, Lift, Swing, Dump, Return). Segmenting into 6 distinct subtasks standardizes execution and prevents compounding geometric errors.

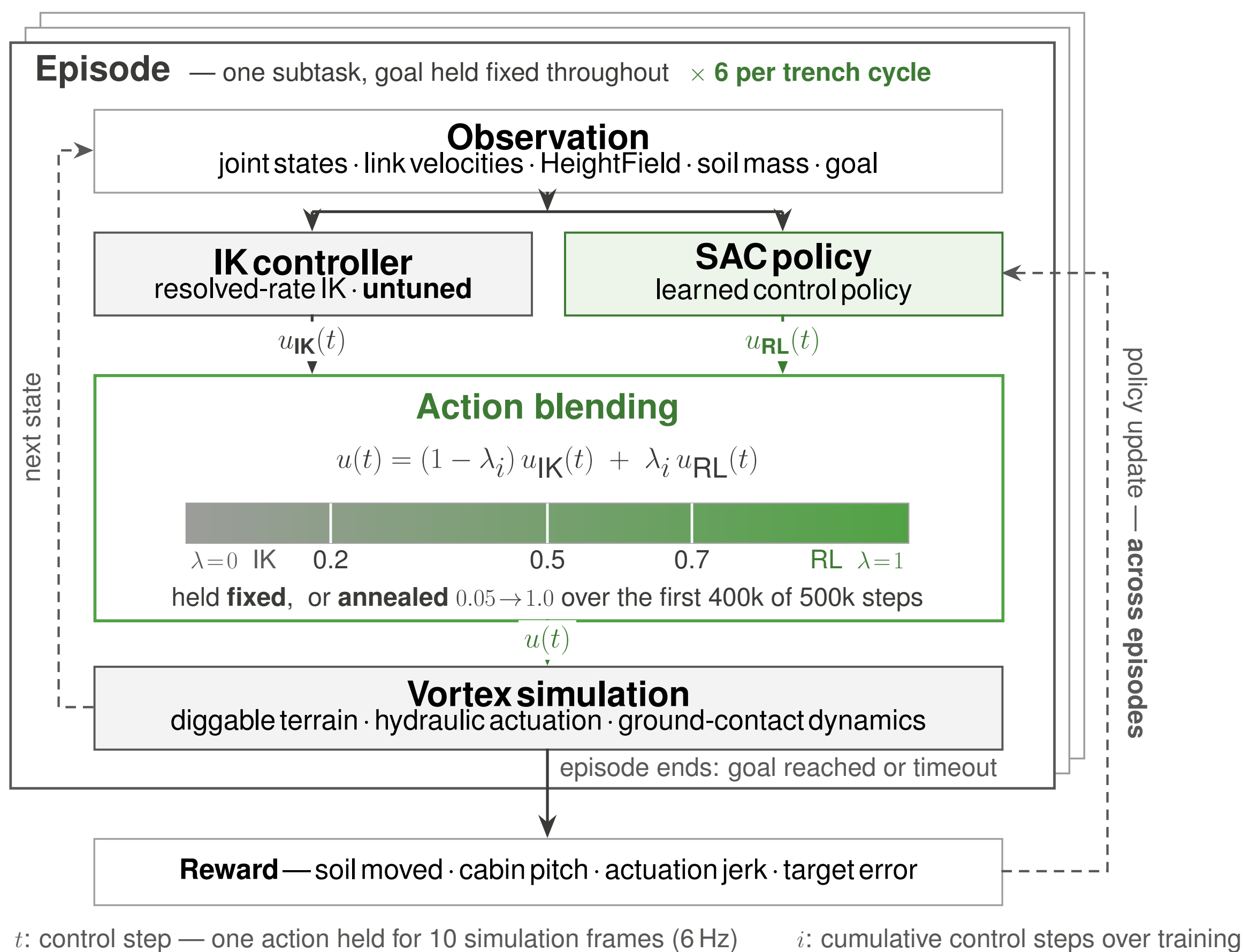


**Fig. 1 The six-step trenching cycle.** Each stage has one checkable goal — a bucket pose and a payload (\$2) — and a run succeeds only when **all six complete in sequence**.

## 2. Learning Methodology

Segmenting the cycle into six short, goal-conditioned episodes keeps credit assignment tractable; within each, a classical IK command and a learned policy are blended into one action [1, 2]. The RL share of authority  $\lambda$  is either held **fixed** for a whole run, or **annealed** from IK toward RL as training proceeds — an **Action Blending Curriculum over Residual Reinforcement Learning (ABC-RRL)**.

**The IK controller is untuned** — a resolved-rate Jacobian controller on a trapezoidal velocity profile, with no per-task tuning. It is guidance for RL to improve on, not a baseline to beat.



**Fig. 2 One goal-conditioned episode.**  $\lambda$  is the RL share of authority: at  $\lambda=0$  the action is the IK command alone, at  $\lambda=1$  the policy has the machine.

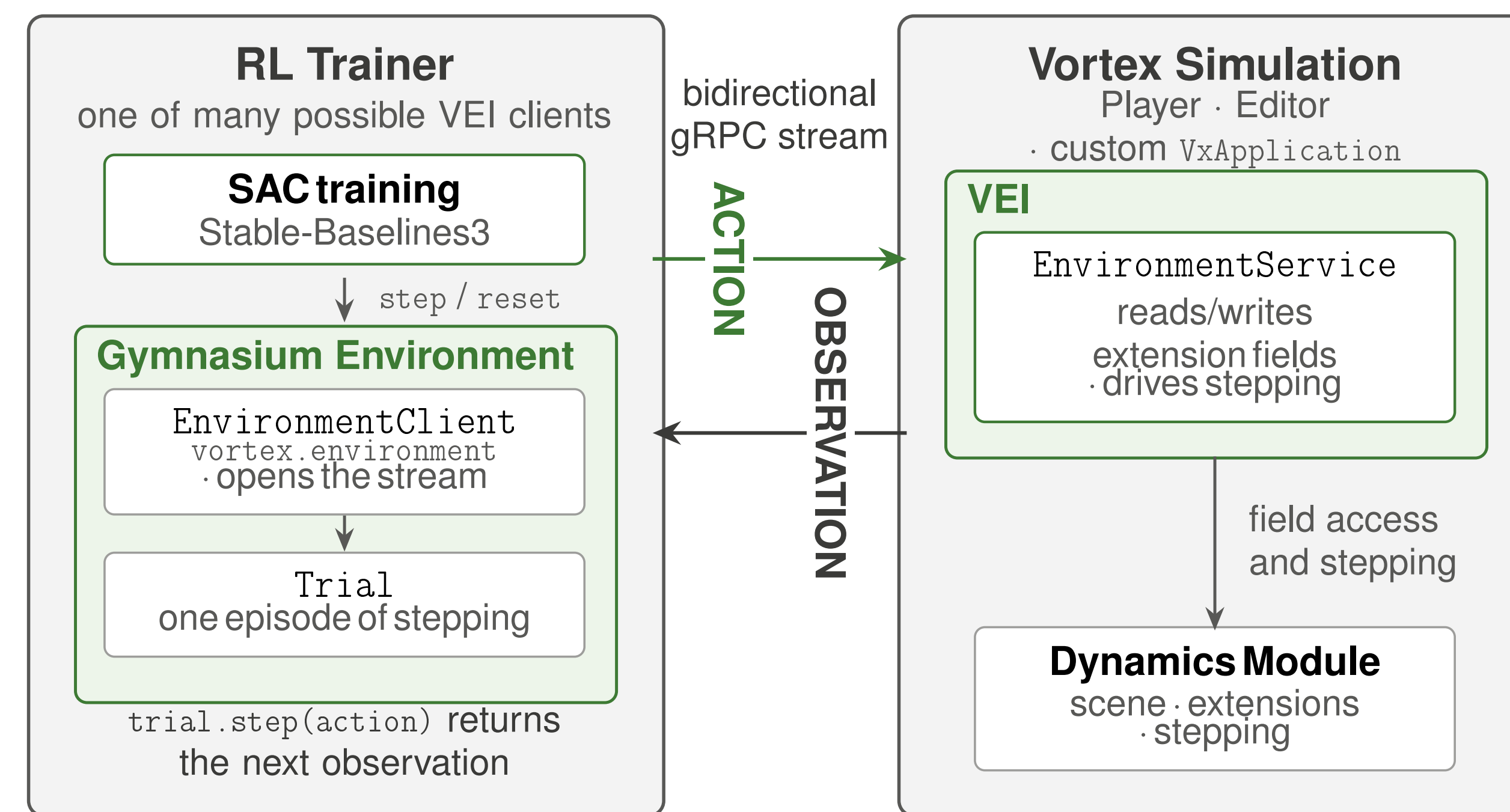
**Goal conditioning.** One policy covers all six subtasks: no subtask IDs are passed, only the dual target vector  $[e_t, m_{\text{target}}]$ .

|                                       |                                                                                                                                                                                                                                              |
|---------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Spatial target error</b>           | $e_t = x_{\text{bucket}}(t) - g_i$ for active goal $g_i$                                                                                                                                                                                     |
| <b>Payload target &amp; tolerance</b> | $3000 \pm 1000$ kg ( <i>Dig</i> ) · $2500 \pm 800$ kg ( <i>Lift</i> ) · $2500 \pm 1000$ kg ( <i>Swing</i> ) · 0 kg elsewhere                                                                                                                 |
| <b>Stage complete when</b>            | bucket within <b>0.5 m</b> and <b>0.1 rad</b> of the goal <i>and</i> payload inside the stage tolerance                                                                                                                                      |
| <b>On subtask failure</b>             | Subtask timeout triggers early termination                                                                                                                                                                                                   |
| <b>Reward</b>                         | $R_t$ rewards step-wise progress in bucket pose ( $R_{\text{bucket}}$ ) and soil payload ( $R_{\text{soil}}$ ) relative to previous step values, penalizes duration to optimize cycle time, and awards a completion bonus on subtask success |

## 3. Platform & Experimental Setup

**The simulation platform is built on Vortex Studio.** The excavator is a production model used in commercial operator training simulators, not a reduced stand-in built for RL. Hydraulic actuation, ground contact and terrain deformation [5] behave as they do for human trainees.

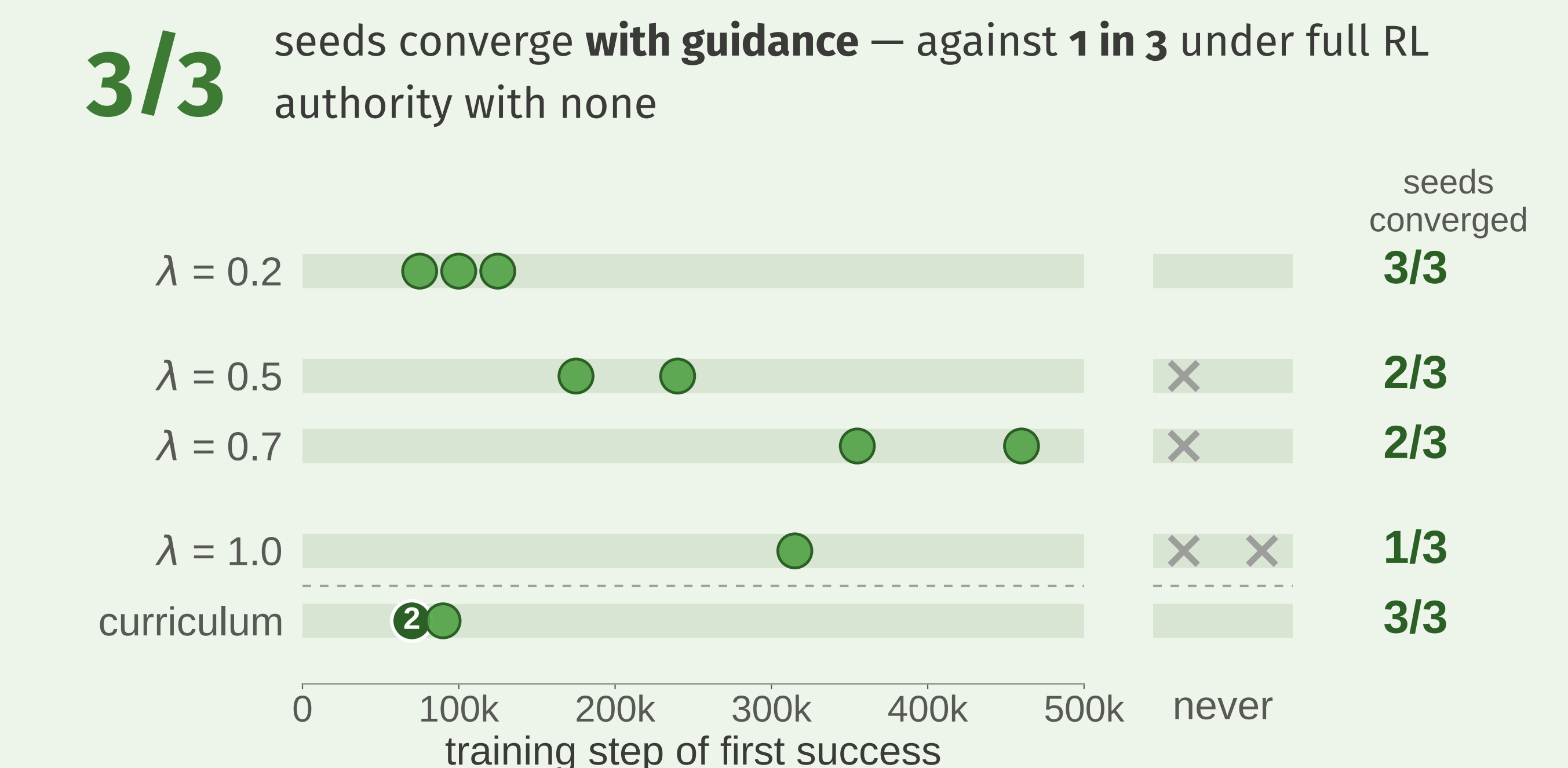
**The Vortex Environment Interface (VEI)** exposes a running Vortex simulation as a standard, synchronous RL environment: one gRPC stream, fields addressed by name, deterministic stepping.



**Fig. 3 The Vortex Environment Interface.** One bidirectional gRPC stream; the RL trainer shown is one possible client, not the only shape.

|                     |                                                                                                                                                                                                                                                                                                                                                                                                                                    |
|---------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Action</b>       | Box(4) — joint commands <i>Swing, Bucket, Stick, Boom</i> , continuous $[-1, 1]$                                                                                                                                                                                                                                                                                                                                                   |
| <b>Observation</b>  | Box(159) — joint positions, linear & angular link velocities; world transforms for turret, arm, boom, chassis; HeightField array (Vortex Earthwork); and the stage's <b>target</b> payload — not the measured one, which feeds only the reward and the success test                                                                                                                                                                |
| <b>Agent</b>        | SAC (Stable-Baselines3 [3]), chosen for stability on continuous action spaces; MlpPolicy $2 \times 256$ , lr $3 \times 10^{-4}$ , batch 256, buffer $10^6$ , $\gamma = 0.99$ , $\tau = 0.005$ , auto entropy; 500k steps per configuration, $\approx 24$ h of wall clock each trained at <b>6 Hz</b> (one action per 10 simulation frames), evaluated at <b>60 Hz</b> — faster control than the training rate improved performance |
| <b>Control rate</b> |                                                                                                                                                                                                                                                                                                                                                                                                                                    |
| <b>Evaluation</b>   | <b>3 seeds</b> per configuration — 15 models, <b>150 episodes</b> each. Spread under $\sim 5\%$ is episode sampling, so no error bars [4]. Two of the 15 are earlier best checkpoints, not 500k                                                                                                                                                                                                                                    |

## 4. Results & Insights



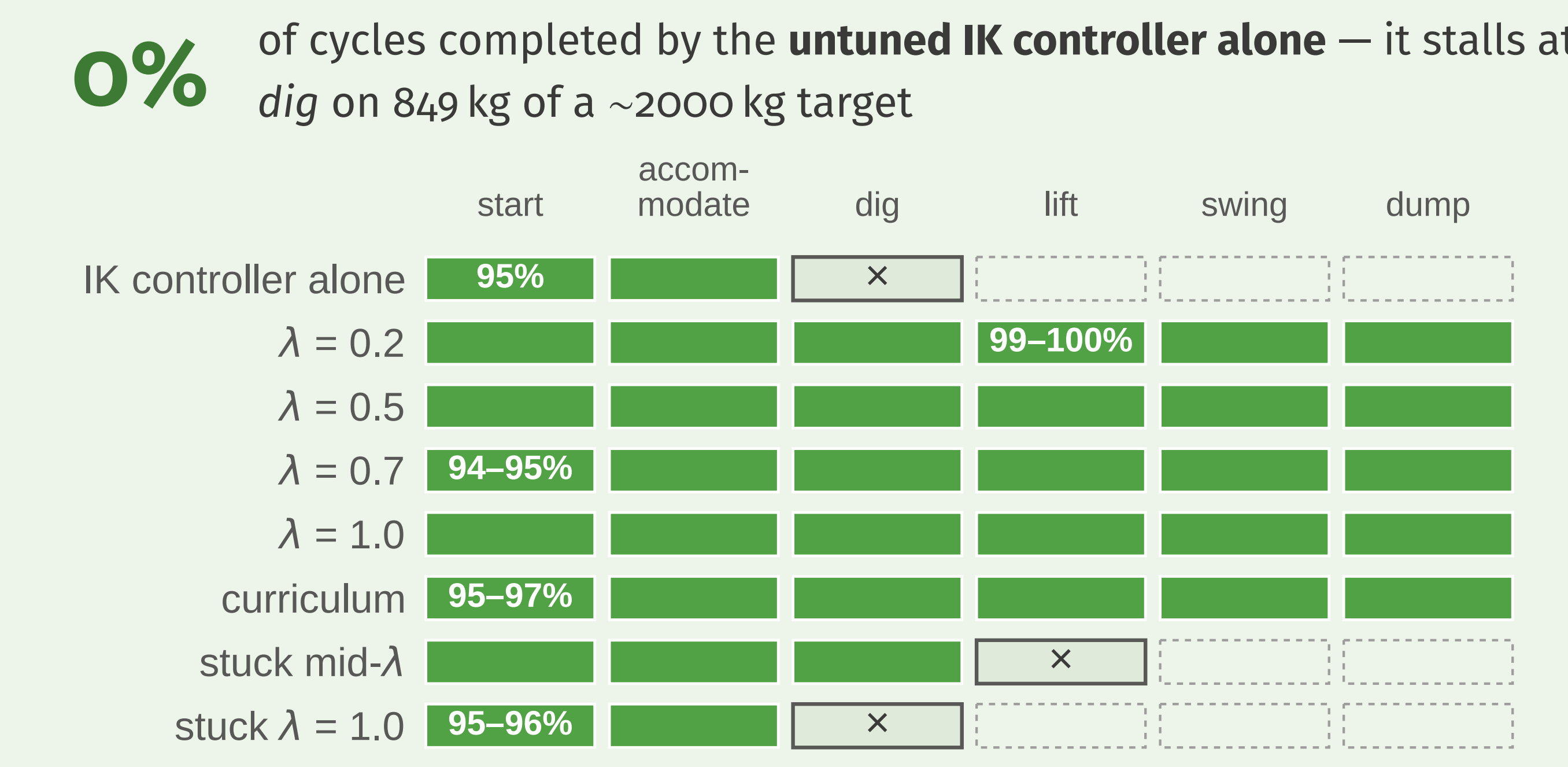
**Fig. 4 Reliability falls as policy authority rises.** One dot per seed, numbered where seeds coincide;  $\lambda$  to scale.  $\times$  = never succeeded inside 500k.

### Conclusions

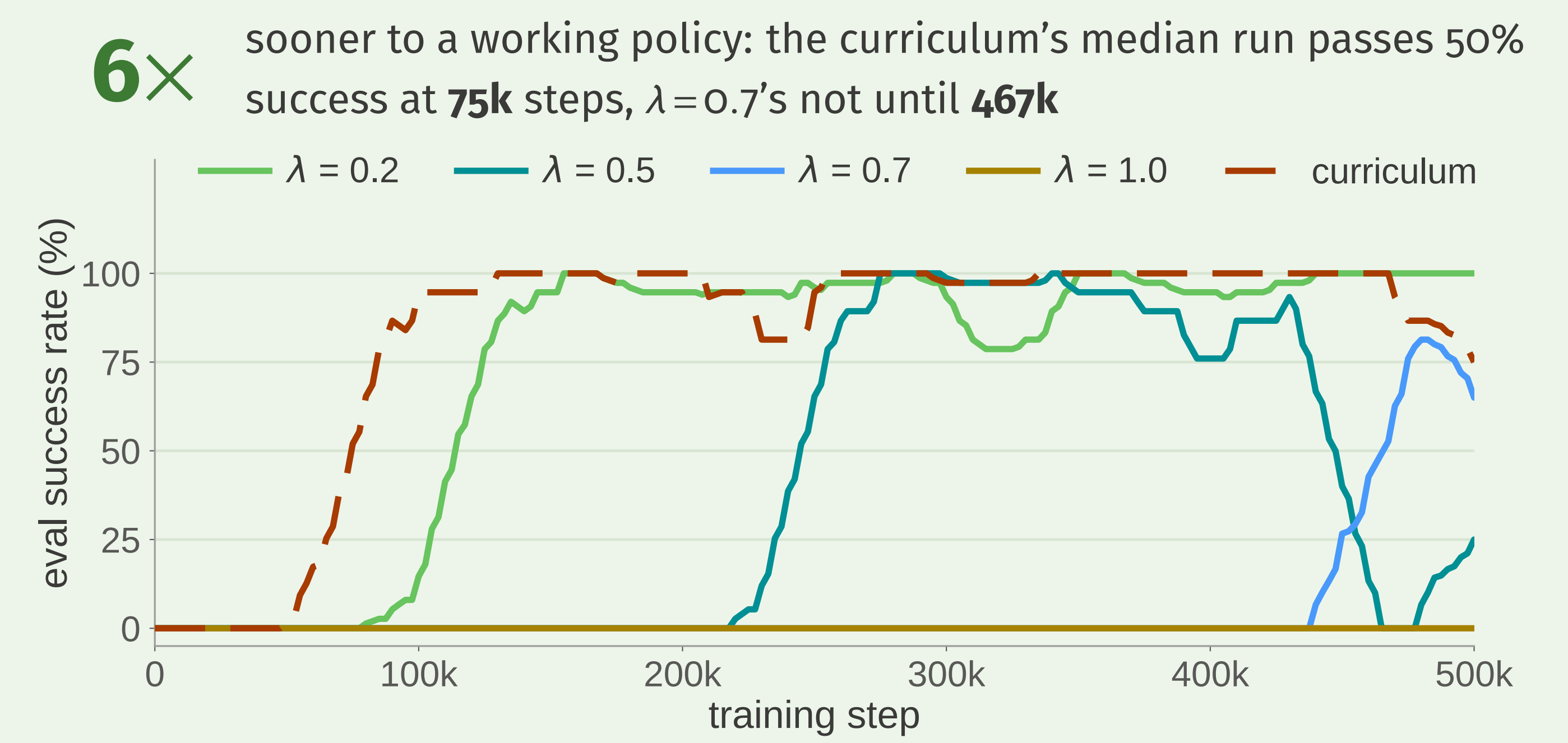
**Neither extreme is reliable.** The untuned IK controller never completes a cycle — it stalls at *dig* on 849 kg. Unguided RL is not impossible, but only 1 seed in 3 converged; the other two dug 24–76 kg, worse than no policy at all.

**Guidance removes the downside, not the ceiling.** The best single model here is unguided ( $\lambda=1$ , 100%). What guidance buys is that *every* seed gets there — 3/3 at  $\lambda=0.2$ , and first success 3–4 $\times$  earlier.

**The curriculum keeps that reliability at full autonomy.** Annealing  $\lambda$  from 0 to 1 gave 3/3 deployable policies at 94.7–97.3% — within sampling noise of the fixed blends, with the IK controller switched out by the end.



**Fig. 5 Each method stops at a different link.** Episodes completing each stage, pooled over the row's seeds; a range where it is not 100%.  $\times$  = stops here, dashed = never reached.



**Fig. 6 Guided runs rise earlier.** Median over all three seeds, 15-point rolling mean. At  $\lambda=1$  it never leaves zero.

### Future Work

Richer reward shaping; other earthmoving tasks, longer horizons and other machines; then sim-to-real transfer. Everything here is simulation, on one soil model and one trench.