

Cast a Wider Net: Coordinated Pass@K Policy Optimization for Code Reasoning

Yilong Li and Suman Banerjee
University of Wisconsin–Madison
{yli758, suman}@wisc.edu

Tong Che*
NVIDIA Research
tongc@nvidia.com

Abstract

While repeated sampling against a verifier is the standard test-time compute strategy for code generation (measured by pass@ K), drawing K independent samples from a single distribution often wastes the compute budget on redundant, near-duplicate reasoning paths. This failure is costly in competitive programming, where many problems admit multiple distinct algorithmic strategies and passing pass@ K requires only one correct attempt. We propose Coordinated Pass@K Policy Optimization (CPPO), which trains a joint planner–solver policy for pass@ K : a planner emits a tuple of $K=4$ alternative high-level methods, and a shared solver attempts one solution per method. We train it with a multiplicative planner reward, $R_{\text{plan}} = J_{\psi} \cdot R_{\text{out}}$, assigning credit only to valid strategy tuples that lead to verifier-confirmed pass@ K success. Across APPS, CodeContests, and LiveCodeBench-v6, CPPO improves pass@4 over direct sampling, planning baselines, planner-only SFT, and pass@ K -oriented RL under the same $K=4$ solver-attempt budget, with statistically significant gains on six of nine model–benchmark cells; for instance, it lifts Qwen3.5-9B LiveCodeBench-v6 from 0.588 (PKPO, the strongest baseline) to 0.728 (+0.14; hierarchical bootstrap, $p < 0.05$).

1 Introduction

Repeated sampling against a verifier is the standard way to allocate test-time compute in code generation (Brown et al., 2024; Snell et al., 2025), with pass@ K as the canonical metric for competitive-programming benchmarks (Chen et al., 2021; Li et al., 2022; Jain et al., 2025). Given a fixed attempt budget, the central question is how to allocate attempts effectively.

Standard pipelines spend this budget on K independent samples from one answer distribution.

Once the model is moderately confident, these samples tend to collapse onto the dominant mode: on a programming problem, this can mean several rollouts that implement the same algorithm with only cosmetic changes—for example, four memoized variants of one dynamic-programming solution—even when the problem admits distinct alternatives. Recent pass@ K training methods such as pass@ K -only RL (Chen et al., 2025) and PKPO (Walder and Karkhanis, 2025), as well as diversity-aware RL methods such as Darling (Li et al., 2025), improve how answer samples are rewarded, but they do not model the K attempts as one coordinated action: the attempts remain separate draws from the answer distribution, so the same mode-dropping failure can persist.

Rather than relying on independent draws, we train the model to generate the K attempts jointly. We focus on code reasoning, where competitive-programming problems often admit multiple algorithmic strategies and any one correct attempt is enough to pass. A planner jointly emits a structured *strategy tuple* of $K=4$ high-level methods, each conditioned on earlier ones, and a shared solver attempts one solution per method. Improving pass@ K therefore becomes a problem of producing complementary alternatives rather than resampling from a single answer distribution. To our knowledge, CPPO is the first RLVR post-training method that optimizes an autoregressive strategy-tuple planner under a verifier-confirmed pass@ K objective. We call this method Coordinated Pass@K Policy Optimization (CPPO; Figure 1); the full training algorithm is in §3.3.¹

Table 1 evaluates CPPO against top baselines under a fixed $K=4$ budget. Full CPPO consistently outperforms Tuple Planner SFT—yielding pass@4 gains of +0.07 to +0.24—which clearly

*Project Lead, Corresponding Author.

¹Here, CPPO denotes *Coordinated Pass@K Policy Optimization*: a joint planner–solver policy trained with the multiplicative planner reward $R_{\text{plan}} = J_{\psi} \cdot R_{\text{out}}$.

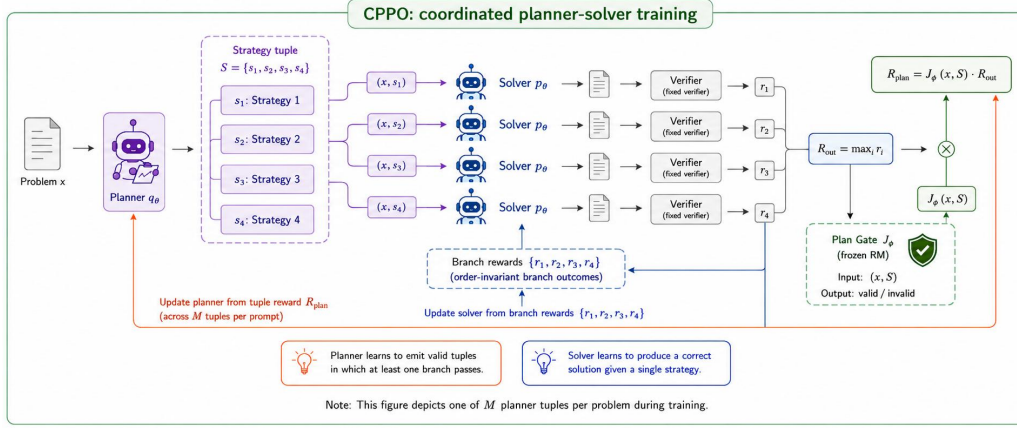


Figure 1: Overview of Coordinated Pass@K Policy Optimization. The planner q_Θ emits a strategy tuple $S = (s_1, \dots, s_K)$; the shared solver p_Θ produces one solution per strategy; a verifier returns per-branch outcomes $r_i \in \{0, 1\}$, and the outcome reward $R_{\text{out}} = \max_i r_i$ scores pass@K success. A frozen reward model $J_\psi(x, S)$ gates plan validity, giving the planner reward $R_{\text{plan}} = J_\psi(x, S) \cdot R_{\text{out}}$, which is nonzero only when the tuple is accepted by J_ψ and at least one of its solver branches passes. Solver tokens are updated from within-tuple outcome advantages; planner tokens are updated from across-tuple advantages of R_{plan} .

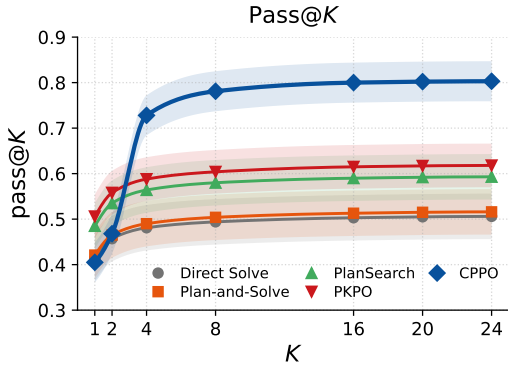


Figure 2: LiveCodeBench-v6 pass@K for Qwen3.5-9B (4B counterpart in Figure 5); all CPPO points use one $K_{\text{tuple}}=4$ planner and pool 4-tuple rollouts for $K>4$ (Table 22).

isolates the value of joint planner–solver RL over SFT initialization alone. The contributions can be summarized:

- **Verifier-trained strategy tuples for pass@K.** We recast pass@K RLVR from independently scored answer samples to an autoregressive planner that emits a joint strategy tuple, separating CPPO from prior pass@K-only RL (Chen et al., 2025) and PKPO (Walder and Karkhanis, 2025).
- **Multiplicative planner reward.** We introduce a multiplicative planner reward,

$$R_{\text{plan}} = J_\psi \cdot R_{\text{out}},$$

that assigns credit to the planner only if the learned validity gate J_ψ accepts the strategy tuple and at least one solver branch succeeds. This

gate acts as a strict validity constraint, primarily suppressing credit for malformed or duplicate plans (§4.5.1), whereas verifier-confirmed solver success (R_{out}) provides the pass@K signal. We optimize this objective end-to-end using split-region GRPO updates for both planner and solver tokens (§3.2).

- **Pass@4 gains at fixed budget.** Across APPS, CodeContests, and LiveCodeBench-v6, CPPO improves pass@4 over the baseline in all nine (size, benchmark) cells of Table 1 and Appendix A under the same $K=4$ solver-attempt budget and verifier, with six cells reaching $p < 0.05$ under a hierarchical problem-and-seed bootstrap; on Qwen3.5-9B LiveCodeBench-v6 the gain over PKPO is +0.14 (0.588 $\rightarrow$ 0.728), and the gains persist at larger attempt budgets (Fig. 2).

2 Background and Related Work

RLVR and optimizer. We work in the Reinforcement Learning with Verifiable Rewards (RLVR) setting popularized by DeepSeek-R1 (Guo et al., 2025): a rule-based verifier $V(x, y) \in \{0, 1\}$ returns whether response y is correct for prompt x , and this binary signal serves as the reward for policy optimization. We use GRPO (Shao et al., 2024) with group-normalized advantages and a clipped, KL-regularized objective (Schulman et al., 2017). CPPO keeps this optimizer fixed and changes only the policy factorization, reward, and token regions

Table 1: Pass@4 results for Qwen3.5-4B and 9B, evaluated with a shared verifier and a constant $K=4$ solver-attempt budget.

Method	Base model	APPS (Hendrycks et al., 2021)	CodeContests (Li et al., 2022)	LiveCodeBench (LCBv6) (Jain et al., 2025)
Direct Solve	Qwen3.5-4B	0.515 ± 0.022	0.094 ± 0.014	0.214 ± 0.013
Direct Solve	Qwen3.5-9B	0.696 ± 0.021	0.175 ± 0.019	0.481 ± 0.014
Plan-and-Solve	Qwen3.5-4B	0.530 ± 0.019	0.098 ± 0.016	0.255 ± 0.012
Plan-and-Solve	Qwen3.5-9B	0.801 ± 0.015	0.171 ± 0.017	0.490 ± 0.020
PlanSearch	Qwen3.5-4B	0.554 ± 0.022	0.128 ± 0.021	0.236 ± 0.012
PlanSearch	Qwen3.5-9B	0.770 ± 0.019	0.168 ± 0.018	0.564 ± 0.019
Pass@ K Training / RLVR	Qwen3.5-4B	0.607 ± 0.017	0.115 ± 0.020	0.316 ± 0.019
Pass@ K Training / RLVR	Qwen3.5-9B	0.762 ± 0.015	0.185 ± 0.021	0.503 ± 0.020
PKPO	Qwen3.5-4B	0.722 ± 0.021	0.156 ± 0.020	0.488 ± 0.015
PKPO	Qwen3.5-9B	0.787 ± 0.020	0.183 ± 0.025	0.588 ± 0.015
UpSkill	Qwen3.5-4B	0.661 ± 0.023	0.121 ± 0.019	0.411 ± 0.016
UpSkill	Qwen3.5-9B	0.762 ± 0.015	0.136 ± 0.017	0.542 ± 0.013
Tuple Planner SFT	Qwen3.5-4B	0.548 ± 0.024	0.147 ± 0.021	0.348 ± 0.018
Tuple Planner SFT	Qwen3.5-9B	0.733 ± 0.016	0.279 ± 0.026	0.514 ± 0.015
CPPO	Qwen3.5-4B	$0.784^\dagger \pm 0.013$	$0.231^\dagger \pm 0.025$	0.505 ± 0.017
CPPO	Qwen3.5-9B	0.806 ± 0.019	$0.396^\dagger \pm 0.028$	$0.728^\dagger \pm 0.016$

Tuple Planner SFT emits all K plans in one autoregressive tuple, $q_\Theta(s_i | x, s_{<i})$; CPPO further trains this planner-solver policy with validity-gated outcome rewards. The matched iid-inference ablation is in Appendix I.

[†] CPPO significantly outperforms the strongest non-CPPO baseline in that cell under a hierarchical bootstrap over problems and training seeds ($p < 0.05$, 1000 resamples; per-seed CIs in Appendix B). Unmarked CPPO gains are positive but not significant.

Baselines. We compare independent sampling, prompt-based planning, inference-time plan search, pass@ K -oriented RL, diversity-oriented RL, and planner-only SFT; details are in §4.3 and Appendix R.

to which advantages are applied; Appendix J gives the standard GRPO equations.

Pass@ K objective. With a fixed attempt budget K per prompt, pass@ K is the indicator that at least one of K attempts is correct,

$$R_{\text{pass@}K}(x, y_{1:K}) = \mathbb{I}\left\{\max_{i=1}^K V(x, y_i) = 1\right\}. \quad (1)$$

The standard repeated-sampling factorization draws the K attempts independently from a single answer distribution $p_\Theta(\cdot | x)$,

$$y_i \sim p_\Theta(\cdot | x), \quad i = 1, \dots, K. \quad (2)$$

Mode coverage under independent sampling.

A simple bound illustrates the limitation of drawing attempts independently. If $p_\Theta(\cdot | x)$ places probability ϵ on a correct-answer region $\mathcal{A}$, then none of K iid samples lands in $\mathcal{A}$ with probability $(1 - \epsilon)^K$; for $\epsilon=0.05$ even $K=16$ leaves a missing-mass probability of ≈ 0.44 . Re-weighting the per-sample reward (pass@ K -only RL, PKPO) reshapes p_Θ —and can raise ϵ —but the K draws stay independent, with no explicit negative dependence or branch-level allocation across strategy regions. CPPO instead introduces a structured joint action: its K branches are sampled jointly through the planner (§3.1) and can be allocated to complementary strategies within a single trajectory.

Repeated sampling and test-time compute. Repeated sampling is the standard way to allocate test-time compute for code generation (Brown

et al., 2024), and compute-optimal sampling can rival substantially larger models (Snell et al., 2025); CPPO targets this same pass@ K regime but changes how the K attempts are generated.²

Pass@ K training and diversity rewards. Recent post-training methods improve repeated sampling by changing how answer samples receive credit. Pass@ K training (Chen et al., 2025) and PKPO (Walder and Karkhanis, 2025) optimize set-level pass@ K rewards over sampled answers, while Darling (Li et al., 2025), DIVER (Hu et al., 2026), and UpSkill (Shah et al., 2026) add diversity pressure through reward shaping, exploration bonuses, or latent-skill regularization. These methods can broaden the answer distribution, but they do not model the K attempts as a single coordinated policy action. CPPO instead treats the K -way set as a structured action: one planner rollout emits a coordinated tuple of strategies, and the solver executes one branch per strategy.

Verifiers and LLM-as-judge. Outcome verifiers (Cobbe et al., 2021) and process reward models (Lightman et al., 2024) score model outputs without ground-truth labels. CPPO’s plan-validity reward model is a small generative verifier in this tradition, with labels obtained via LLM-as-judge (Zheng et al., 2023).

Planning for reasoning. Planning methods add structure before generation rather than after scor-

²Results in this version use an updated training-data setup for the planner (§4.2); the relative ordering of methods in Table 1 and our conclusions are unchanged.

ing. Plan-and-solve prompting (Wang et al., 2023a), Tree of Thoughts (Yao et al., 2023), and PlanSearch (Wang et al., 2024) use natural-language plans or search procedures at inference time over a frozen model. They show that explicit plans can help reasoning, but they do not train a planner under the pass@ K objective. CPPO brings planning into post-training: the planner learns to emit a single autoregressive strategy tuple whose later items condition on earlier ones, and receives validity-gated credit only when the tuple leads to verifier-confirmed pass@ K success.

3 Coordinated Pass@ K Policy Optimization

CPPO changes the policy factorization: instead of drawing K iid answers, it samples one tuple of $K=4$ algorithmic strategies. This factorization splits credit assignment: the solver must learn which strategy-conditioned branch succeeds, and the planner must learn which strategy tuple makes at least one branch succeed. CPPO therefore assigns branch-level outcome rewards to solver tokens and a tuple-level reward to planner tokens. Because tuple-level rewards are sparse and can be gamed by malformed plans, we gate planner credit with a narrow plan-validity reward model and train the system through a staged procedure. Table 2 contrasts CPPO with the closest planning and pass@ K post-training baselines along the design axes that follow.

Table 2: Comparison of CPPO with the closest planning and pass@ K post-training baselines along four axes: output granularity, search strategy, intervention stage, and RLVR training.

Method	Granularity	Search strategy	Stage	RLVR ready
Plan-and-Solve	answer	iid	inference	–
PlanSearch	plan features	enumerate	inference	–
Pass@ K -only RL	answer	iid	post-training	✓
PKPO	answer	iid	post-training	✓
CPPO (ours)	strategy	per-item	both	✓

3.1 A Joint Strategy-Tuple Policy

Standard pass@ K samples K answers independently from one distribution. CPPO instead samples one structured object: a tuple of K strategy sketches generated autoregressively, so each method conditions on its predecessors and avoids redundant branches. The planner emits the tuple

$$S = (s_1, \dots, s_K), \quad (3)$$

with factorization

$$q_{\Theta}(S | x) = \prod_{i=1}^K q_{\Theta}(s_i | x, s_{<i}). \quad (4)$$

A shared solver produces one answer per method,

$$y_i \sim p_{\Theta}(\cdot | x, s_i), \quad i = 1, \dots, K, \quad (5)$$

yielding the full trajectory

$$\tau = (S, y_{1:K}) \quad (6)$$

with probability

$$\pi_{\Theta}(\tau | x) = q_{\Theta}(S | x) \prod_{i=1}^K p_{\Theta}(y_i | x, s_i). \quad (7)$$

In the simplest implementation, a single backbone with shared parameters Θ induces both q_{Θ} (planner factor) and p_{Θ} (solver factor) from two prompt modes: the PLAN prompt produces a concise, task-focused strategy tuple (planner tokens), and the SOLVE prompt produces one answer conditioned on a single strategy (solver tokens). Planner and solver losses then apply to these disjoint token regions (exact prompts in Appendix D).

3.2 Coordinated Rewards and Credit Assignment

CPPO assigns credit at two granularities. Solver tokens receive branch-level verifier feedback: for branch i , $r_i^{\text{out}} = V(x, y_i)$, and the tuple-level outcome reward is the pass@ K indicator

$$R_{\text{out}}(x, y_{1:K}) = \mathbb{I}\left\{\max_i V(x, y_i) = 1\right\}. \quad (8)$$

This outcome signal is appropriate for solver branches, but it does not tell the planner whether the visible tuple is well formed, non-duplicative, or free of answer leakage. We therefore gate planner credit with a narrow plan-validity model $J_{\psi}(x, S)$. The gate checks that the tuple parses into K methods, avoids literal duplicates, contains no final answer or completed implementation, specifies actionable approaches, and stays on topic. It does not predict whether a plan will solve the task. CPPO uses its binary decision

$$J_{\psi}(x, S) = \mathbb{I}\{p_{\psi}(\text{Pass} | x, S) \geq \tau\} \in \{0, 1\}, \quad (9)$$

with τ chosen for high recall on validation data (§4.1). Since the gate learns only this minimal contract, a small reward model suffices; training details

are given in Stage 2 of §3.3, and §4.4 validates its use as a gate rather than a quality scorer. The planner reward combines validity and outcome:

$$R_{\text{plan}}(x, S, y_{1:K}) = J_{\psi}(x, S) \cdot R_{\text{out}}(x, y_{1:K}). \quad (10)$$

Thus planner credit is given only to tuples that are valid and yield at least one passing strategy-conditioned solution: R_{out} rejects valid but unsuccessful plans, while J_{ψ} rejects successes obtained through malformed, duplicated plans or plans that leak the answer.

Split-region advantages. The two token regions carry different credit granularity. Solver advantages compare the K branches within one tuple τ ,

$$a_i(\tau, x) = \frac{r_i^{\text{out}} - \text{mean}(r_{1:K}^{\text{out}})}{\text{std}(r_{1:K}^{\text{out}}) + \epsilon}, \quad (11)$$

and update answer tokens through $\log p_{\Theta}(y_i | x, s_i)$. Planner advantages normalize R_{plan} across M tuples sampled for the same prompt,

$$A_{\text{plan}}^{(m)} = \frac{R_{\text{plan}}^{(m)} - \text{mean}(R_{\text{plan}}^{(1:M)})}{\text{std}(R_{\text{plan}}^{(1:M)}) + \epsilon}, \quad (12)$$

and update planner tokens through $\log q_{\Theta}(S^{(m)} | x)$. The two advantages enter the standard clipped, KL-regularized GRPO objective of (18):

$$\mathcal{L} = \sum_{m=1}^M \sum_{i=1}^K \mathcal{L}_{\text{solve}}^{(m,i)}(a_i^{(m)}) + \sum_{m=1}^M \mathcal{L}_{\text{plan}}^{(m)}(A_{\text{plan}}^{(m)}), \quad (13)$$

where $\mathcal{L}_{\text{solve}}$ and $\mathcal{L}_{\text{plan}}$ denote the negative GRPO objective on solver and planner token spans. In the shared-backbone implementation of §3.1, the planner and solver losses apply to disjoint token spans (selected by the PLAN and SOLVE prompts) but update the same parameters Θ .

3.3 Making the Sparse Joint Reward Trainable

The planner reward $R_{\text{plan}} = J_{\psi} \cdot R_{\text{out}}$ is nonzero only when the planner emits a valid tuple and at least one solver branch passes. Since clean K -method tuples and successful rollouts are rare at initialization, CPPO first raises reward density through planner SFT, validity-gate training, RM-guided warm-up, and a reward-density audit.

Stage 1: Planner SFT. We initialize the planner by maximum likelihood on strict- K gold tuples—self-generated plan samples filtered by a Qwen3.5-9B judge (Appendix E.4). All experiments use the same $K=4$ planner; the K sweep in §4.5.1 obtains smaller or larger budgets by truncating or pooling planner rollouts at inference time, without retraining. The CPPO reward is set-level, and the SFT data does not assign semantic roles to tuple positions; to confirm that truncating to the first $K < 4$ outputs introduces no material positional artifact, an order-randomized control in Appendix I reports first- K , last- K , and random- K selection within seed variance of each other.

Stage 2: Validity gate. We train p_{ψ} on a balanced CodeContests tuple pool of LLM-judged (Zheng et al., 2023) strict-four positives and semantically invalid but superficially well-formed negatives, and accept a checkpoint only when pre-registered validation metrics pass. During CPPO, J_{ψ} is frozen within each phase and refreshed between phases on a fresh batch of planner outputs; exact splits, thresholds, and the refresh schedule are in Appendix F.

Stage 3: RM-guided warm-up. We precede joint training with a planner-only stage under $R_{\text{warm}}(x, S) = J_{\psi}(x, S)$, applied to planner tokens through the GRPO objective of (18), until the fraction of plans accepted by J_{ψ} plateaus.

Stage 4: Reward-density audit. We run a forward-only evaluation pass and require a nonzero frozen-solver pass@ K rate from the warmed planner together with a nonzero R_{plan} density across sampled rollouts; if either fails, we return to Stage 1 or Stage 2 rather than launch CPPO.

Stage 5: Joint CPPO update. The inner loop of joint training is summarized in Algorithm 1: M planner tuples per prompt, K branches each, within-tuple solver advantages and across-tuple planner advantages feed the split-region GRPO loss (13). The full pipeline pseudocode is in Appendix K.

Algorithm 1 One CPPO update step (Stage 5).

Require: prompt x , M planner tuples per prompt, K branches per tuple

- 1: Sample $S^{(m)} \sim q_{\Theta}(\cdot | x)$ for $m = 1, \dots, M$.
- 2: Score validity $J_{\psi}(x, S^{(m)})$ with the frozen RM.
- 3: Solve $y_i^{(m)} \sim p_{\Theta}(\cdot | x, s_i^{(m)})$ for every branch.
- 4: Verify $r_i^{(m)} = V(x, y_i^{(m)})$ and form $R_{\text{out}}^{(m)}, R_{\text{plan}}^{(m)} = J_{\psi}(x, S^{(m)}) \cdot R_{\text{out}}^{(m)}$.
- 5: Compute within-tuple solver advantages $a_i^{(m)}$ from (11).
- 6: Compute across-tuple planner advantages $A_{\text{plan}}^{(m)}$ from (12).
- 7: Update Θ with the split-region GRPO loss (13).

4 Experimental Setup

We instantiate CPPO on competitive programming, where the verifier V executes each candidate program against the problem’s tests inside a sandboxed, time-limited runner (§3.2).

4.1 Models

Unless otherwise stated, the planner and solver share the same backbone and are switched by the PLAN and SOLVE prompt modes of §3.1; the two token regions are trained jointly in full CPPO. The main experiments evaluate Qwen3.5- $\{2B, 4B, 9B\}$ (Qwen Team, 2026); we additionally verify generality on Gemma 4 (Google DeepMind, 2026) in §4.6. The plan-validity reward model is a separate smaller generative verifier (Qwen3.5-0.8B, §3.2) frozen during warm-up and CPPO. We choose its decision threshold τ in (9) on a validation set to obtain a high-recall validity gate; the main experiments use $\tau=0.17$ and $M=8$ planner tuples per prompt for the across-tuple advantage of (12). Trainable parameters use full precision: at CPPO-scale learning rates the AdamW (Loshchilov and Hutter, 2019) update falls below the bf16 mantissa resolution, rendering half-precision optimization numerically ineffective. The frozen reference and reward models stay in half precision.

4.2 Training Data

All training draws from CodeContests (Li et al., 2022): planner SFT, reward-model data, the planner warm-up, and the joint CPPO rollout pool all use CodeContests train, whose problems admit enough alternative solution strategies to support tuple-level coordination. APPS (Hendrycks et al., 2021) and LiveCodeBench-v6 (LCBv6) (Jain et al., 2025) are *evaluation-only*: neither ever enters any training-side pool—no planner SFT, no LLM-judge label, no reward-model train/val, no warm-up, no CPPO rollout, no hyperparameter sweep, no qualitative example—so both are held-out cross-corpus transfer benchmarks, while CodeContests valid

Table 3: Stage-wise data sources. Held-out evaluation sets are disjoint from all training-side pools under the decontamination protocol in Appendix E.

Stage	Role	Source split
Planner SFT	train	CodeContests train
Reward model	train/val	CodeContests train, filtered
Planner warm-up	train	CodeContests train, filtered
Joint CPPO	train	CodeContests train
APPS	eval	APPS test
CodeContests	eval	CodeContests valid
LCBv6	eval	LCBv6 held-out

measures in-domain pass@ K . Table 3 lists the source split used by each stage; problem counts and the full per-component exposure audit are in Appendix E (Table 10).

Decontamination protocol. We deduplicate each corpus before splitting. If a training-side item matches a held-out evaluation problem, we remove it from training and keep the evaluation split fixed. We match problems by official identifiers when available (Codeforces contest/index, APPS id, LCBv6 id), and otherwise by a canonicalized statement hash with fuzzy n -gram matching (Appendix E.3). After filtering, all APPS–CodeContests, APPS–LCBv6, and CodeContests–LCBv6 train–evaluation pairs have zero prompt-hash overlap.

Candidate selection funnel. Planner SFT uses an over-generate–then-filter pipeline over self-generated plan samples; the candidate-generation prompt is in Appendix D and the rest of the construction is in Appendix E.4.

4.3 Evaluation

Metric and benchmarks. Our primary metric is pass@ K under a fixed verifier budget on three competitive-programming benchmarks: APPS (Hendrycks et al., 2021), CodeContests (Li et al., 2022), and LiveCodeBench-v6 (LCBv6) (Jain et al., 2025). A CPPO rollout consumes one planner call and K solver calls, so we report both verifier/solver-attempt pass@ K (matching prior work) and token-normalized pass@ K (Table 4), which accounts for the planner-call overhead. The full APPS pass@ K sweep across base-lines and model sizes is in Appendix P.

Verifier protocol. Candidate programs are extracted with the same parser for every method and executed in an isolated sandbox. Timeout, memory overflow, compilation error, runtime error, empty

Table 4: Decoded-token-normalized pass@4 for Qwen3.5-4B at $K=4$ on the held-out APPS and LiveCodeBench-v6 evaluation sets (failed and empty branches included). Prompt tokens are excluded; full accounting in Appendix M.

Dataset	Method	pass@4	Decoded toks.	pass@4 / 10k toks.
APPS	Direct Solve	0.515	6087	0.846
APPS	Plan-and-Solve	0.530	7326	0.723
APPS	PlanSearch	0.554	5535	1.001
APPS	CPPO	0.784	4728	1.658
LCBv6	Direct Solve	0.214	1731	1.236
LCBv6	Plan-and-Solve	0.255	3975	0.642
LCBv6	PlanSearch	0.236	2137	1.104
LCBv6	PKPO	0.488	2137	2.284
LCBv6	UpSkill	0.411	2218	1.853
LCBv6	CPPO	0.505	1556	3.246

extraction, and wrong answer are all scored as $V(x, y)=0$. We use the official benchmark test cases when available, and never expose evaluation tests during training. Sandbox configuration (time-out, memory cap, network and filesystem isolation) is reported in Appendix E.5.

Baselines. CPPO is compared against (i) independent answer sampling at the same K ; (ii) independent sketch-then-solve, where each sketch is sampled without seeing the others; (iii) self-consistency and other inference-time aggregation (Wang et al., 2023b); (iv) planning and search baselines—Plan-and-Solve (Wang et al., 2023a) and PlanSearch (Wang et al., 2024); (v) pass@ K training without a planner reward tied to tuple validity and solver success (Chen et al., 2025; Walder and Karkhanis, 2025); and (vi) ablations that remove the plan-validity reward or the outcome reward from the planner.

Token accounting. Table 4 reports token-normalized pass@4 on the held-out APPS and LiveCodeBench-v6 evaluation sets (Table 10). On these sets, the concise PLAN contract yields shorter strategy-conditioned solver completions, more than offsetting the planner generation in decoded-token accounting (planner–solver breakdown in Appendix M).

Variance and significance. For training-based methods—Tuple Planner SFT, RM warm-up, Full CPPO, Pass@ K Training / RLVR, PKPO, and UpSkill—we report $\text{mean} \pm \text{std}$ over three independent training seeds; each seed re-runs the relevant parameter-updating stages end-to-end under the same data splits and hyperparameters. For inference-only baselines that we run ourselves (Direct Solve, Plan-and-Solve, PlanSearch), variation

is over evaluation problems only. Pairwise significance ($\dagger$) uses a *hierarchical bootstrap* that resamples both evaluation problems and training seeds: each of 1000 resamples draws problems with replacement and, for each trained method, draws one of its three seeds with replacement, so the reported p -values reflect evaluation noise *and* training-seed variability rather than being conditional on a fixed checkpoint. With only three seeds this seed-level test is deliberately conservative; Appendix B additionally reports the per-seed CPPO-minus-strongest-baseline difference, its mean, and a 95% confidence interval over seeds. Externally reported numbers, such as the LCBv6 Direct Solve cells of Gemma 4 in Table 21, use the uncertainty provided by the original source when available; otherwise we report the point estimate and exclude those cells from significance tests.

4.4 Validity gate diagnostics

The plan-validity RM fits held-out LLM-judge labels well (AUC 0.971; Figure 3a) but only weakly predicts frozen-solver pass@ K outcomes (AUC 0.572). At $\tau=0.17$, however, accepted tuples pass more often than rejected ones (35% vs. 15%), and only accepted-and-solved tuples receive nonzero R_{plan} (Figure 3b). Thus J_ψ serves as a validity filter: it suppresses malformed or duplicated tuples and plans that leak the answer, as shown by the duplicate-rate jump under $-J_\psi$ in Table 6 (0.08 $\rightarrow$ 0.32). It is too noisy to rank surviving plans by solver utility, so we use it only to gate training rollouts and leave R_{out} as the quality signal.

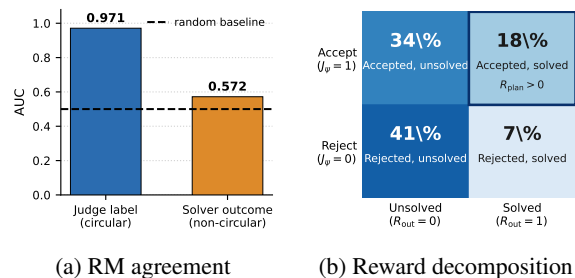


Figure 3: Reward-model diagnostics. (a) The RM matches held-out judge labels but weakly predicts frozen-solver outcomes, supporting its use as a validity gate rather than a plan-quality scorer. (b) Joint distribution of validity decisions and solver outcomes during CPPO rollouts; shading encodes rollout frequency, and only the *accepted, solved* cell yields nonzero R_{plan} .

4.5 Ablation

Table 5 isolates the contribution of each stage of the CPPO pipeline. Each row adds one stage on top of the previous one: Direct Solve, planner-only SFT, planner SFT followed by RM-guided warm-up, and the full CPPO run.

The ablation separates three effects. First, planner-only SFT is not a complete explanation for CPPO’s gains: it improves over Direct Solve on every cell, but by a margin far smaller than full CPPO. Second, RM-guided warm-up is most useful on CodeContests and is modest or mixed elsewhere, consistent with its role as a validity constraint rather than a task-quality signal. Third, the full CPPO update improves over the warm-up stage in every model–dataset cell, indicating that the main gains come from joint planner–solver RL with verifier-confirmed outcome credit.

4.5.1 Reward-component and rollout ablations

Table 6 isolates CPPO’s two main design choices: the validity gate and across-tuple planner normalization. Removing J_ψ quadruples the duplicate-plan rate while leaving pass@4 close to full CPPO, showing that the gate mainly suppresses invalid planner credit rather than scoring task quality. The J_ψ -only variant removes outcome feedback, while $M=1$ removes across-tuple contrast by replacing (12) with a single-sample REINFORCE update. Increasing M improves the across-tuple advantage with diminishing returns by $M=8$ (we did not test beyond). The bottom block sweeps $K \in \{2, 4, 8\}$ using the same $K_{\text{tuple}}=4$ planner from Table 1, with truncation for $K<4$ and pooling for $K>4$. Appendix P gives the full larger-budget decomposition, and Appendix I reports the matched joint-vs.-iid plan-sampling ablation.

4.6 Majority-vote consistency (maj@4)

Pass@ K rewards *any* successful branch, so a method that spreads probability mass across strategies could in principle inflate pass@ K while leaving the modal answer unchanged—or worse, diluting it on tasks with a single canonical output. We test this directly on the LiveCodeBench output-prediction subsets (code-execution and test-output prediction), where each problem has exactly one correct answer and maj@ K is well-defined: for each problem we draw $K=4$ samples and score the most frequent output against the canonical label,

breaking ties by sample order. Table 7 reports the results for Qwen3.5-2B and Qwen3.5-4B.

CPPO wins both sub-tasks at both sizes: on Qwen3.5-2B, average maj@4 of 0.360 versus 0.205 for Direct Solve and 0.295 for PKPO; on Qwen3.5-4B, 0.430 versus 0.335 and 0.410 respectively. Coordinated planning tightens rather than diffuses the modal answer: diversity across high-level strategies is consistent with sharper outputs once a strategy is chosen.

Generality across base models. The same trend holds on Gemma 4 E2B and E4B (Google DeepMind, 2026) under an identical pipeline (Appendix O).

5 Conclusion

We presented CPPO, a planner–solver method that allocates a fixed pass@ K budget to complementary solution strategies rather than repeated samples from one answer distribution. By tying planner credit to both plan validity and verifier-confirmed solver success, CPPO encourages its K attempts to cover different algorithmic strategies rather than paraphrases of a single approach. Across APPS, CodeContests, and LiveCodeBench-v6, it improves pass@4 over direct sampling, planning baselines, planner-only SFT, and pass@ K -oriented RL on Qwen3.5- $\{2B, 4B, 9B\}$, with the same trend carrying over to Gemma 4 E2B and E4B. Future improvements in pass@ K may depend less on sampling more answers and more on coordinating which strategies those answers pursue.

Limitations

Scope. Our claims are restricted to competitive-programming benchmarks, where problems admit several distinct algorithmic strategies. Math benchmarks such as MATH or AIME usually have a single canonical solution path and are outside our scope—there, CPPO may reduce to a more expensive form of plan-and-solve.

Dependencies. The planner reward $R_{\text{plan}} = J_\psi \cdot R_{\text{out}}$ is sparse by construction, so CPPO requires the warm-up and audit gate to train at all. It also inherits the biases of the offline LLM judge and uses a shared planner–solver backbone; splitting them is left to future work.

Ethics Statement

CPPO trains code policies with execution-based rewards. The verifier checks functional correctness

Table 5: Ablation over the CPPO pipeline on APPS (Hendrycks et al., 2021), CodeContests (Li et al., 2022), and LiveCodeBench-v6 (LCBv6) (Jain et al., 2025). All numbers are pass@4.

Method	Base model	APPS	CodeContests	LiveCodeBench (LCBv6)
Direct Solve	Qwen3.5-2B	0.420 $\pm$ 0.017	0.040 $\pm$ 0.012	0.161 $\pm$ 0.014
Direct Solve	Qwen3.5-4B	0.515 $\pm$ 0.022	0.094 $\pm$ 0.014	0.214 $\pm$ 0.013
Direct Solve	Qwen3.5-9B	0.696 $\pm$ 0.021	0.175 $\pm$ 0.019	0.481 $\pm$ 0.014
Tuple Planner SFT	Qwen3.5-2B	0.453 $\pm$ 0.020	0.087 $\pm$ 0.020	0.232 $\pm$ 0.017
Tuple Planner SFT	Qwen3.5-4B	0.548 $\pm$ 0.024	0.147 $\pm$ 0.021	0.348 $\pm$ 0.018
Tuple Planner SFT	Qwen3.5-9B	0.733 $\pm$ 0.016	0.279 $\pm$ 0.026	0.514 $\pm$ 0.015
Tuple Planner SFT + RM warm-up	Qwen3.5-2B	0.461 $\pm$ 0.024	0.105 $\pm$ 0.018	0.231 $\pm$ 0.015
Tuple Planner SFT + RM warm-up	Qwen3.5-4B	0.564 $\pm$ 0.023	0.204 $\pm$ 0.021	0.384 $\pm$ 0.013
Tuple Planner SFT + RM warm-up	Qwen3.5-9B	0.744 $\pm$ 0.019	0.313 $\pm$ 0.021	0.541 $\pm$ 0.018
CPPO	Qwen3.5-2B	0.536 $\pm$ 0.017	0.201 $\pm$ 0.021	0.418 $\pm$ 0.017
CPPO	Qwen3.5-4B	0.784 $\pm$ 0.013	0.231 $\pm$ 0.025	0.505 $\pm$ 0.017
CPPO	Qwen3.5-9B	0.806 $\pm$ 0.019	0.396 $\pm$ 0.030	0.728 $\pm$ 0.016

Table 6: Planner-reward ablations and pass@ K sweep on APPS / Qwen3.5-2B. The top block contains two sub-experiments: the $-J_\psi$, J_ψ -only rows vary the reward (relative to Full CPPO, $R_{\text{plan}} = J_\psi \cdot R_{\text{out}}$, $M=8$), while the $M \in \{1, 2, 4\}$ rows hold the full $R_{\text{plan}} = J_\psi \cdot R_{\text{out}}$ reward fixed and vary only the across-tuple sample count M . The bottom block varies the solver-attempt budget K . Duplicate rate measures duplicate-plan tuples; lower is better. Uncertainties follow the convention of §4.3.

Planner-reward variants ($M=8$)		
Variant	pass@4 $\uparrow$	Dup. rate $\downarrow$
Full CPPO ($J_\psi R_{\text{out}}$)	0.536 $\pm$ 0.017	0.08 $\pm$ 0.030
$-J_\psi$ gate (R_{out})	0.522 $\pm$ 0.023	0.32 $\pm$ 0.044
J_ψ -only	0.503 $\pm$ 0.024	0.14 $\pm$ 0.032
Across-tuple sample count M ($R_{\text{plan}} = J_\psi \cdot R_{\text{out}}$)		
$M=4$	0.535 $\pm$ 0.023	0.10 $\pm$ 0.032
$M=2$	0.515 $\pm$ 0.022	0.12 $\pm$ 0.034
$M=1$	0.486 $\pm$ 0.022	0.16 $\pm$ 0.035
Attempt-budget sweep		
K	Direct Solve	CPPO
2	0.399 $\pm$ 0.022	0.518 $\pm$ 0.021
4	0.420 $\pm$ 0.017	0.536 $\pm$ 0.017
8	0.432 $\pm$ 0.024	0.561 $\pm$ 0.024

only, not security or robustness, so downstream use requires review. All generated code runs in a sandboxed, time-bounded runner. We use only public APPS, CodeContests, and LiveCodeBench-v6 data under their original licenses and no PII. The plan-validity reward model may inherit biases from the offline LLM judge, so we validate it on balanced held-out data before use.

Code and artifact release. Upon acceptance, we will release code, RM checkpoints, judge prompts, and sandbox configuration, but not benchmark content.

References

Bradley Brown, Jordan Juravsky, Ryan Ehrlich, Ronald Clark, Quoc V. Le, Christopher Ré, and Azalia Mirhoseini. 2024. Large language monkeys: Scaling infer-

ence compute with repeated sampling. *arXiv preprint arXiv:2407.21787*.

Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, and 1 others. 2021. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*.

Zhipeng Chen, Xiaobo Qin, Youbin Wu, Yue Ling, Qinghao Ye, Wayne Xin Zhao, and Guang Shi. 2025. Pass@k training for adaptively balancing exploration and exploitation of large reasoning models. *arXiv preprint arXiv:2508.10751*.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.

Google DeepMind. 2026. Gemma 4: Open lightweight models. Official model cards, <https://ai.google.dev/gemma>. Gemma 4 E2B and E4B. LiveCodeBench-v6 pass@4 values used in this paper (44.0% for E2B, 52.0% for E4B) reproduced from the official model-card evaluation tables. Accessed 2026-05-26.

Daya Guo, Dejian Yang, Haowei Zhang, and 1 others. 2025. *DeepSeek-R1: Incentivizing reasoning capability in LLMs via reinforcement learning*. *Nature*, 645:633–638.

Dan Hendrycks, Steven Basart, Saurav Kadavath, Mantas Mazeika, Akul Arora, Ethan Guo, Collin Burns, Samir Puranik, Horace He, Dawn Song, and Jacob Steinhardt. 2021. Measuring coding challenge competence with APPS. In *Advances in Neural Information Processing Systems (NeurIPS) Datasets and Benchmarks Track*.

Zican Hu, Shilin Zhang, Yafu Li, Jianhao Yan, Xuyang Hu, Leyang Cui, Xiaoye Qu, Chunlin Chen, Yu Cheng, and Zhi Wang. 2026. Diversity-incentivized exploration for versatile reasoning. In

Table 7: Majority-vote accuracy (maj@ K) at $K=4$ on the LiveCodeBench output-prediction tasks—code execution and test-output prediction—for Qwen3.5-2B and Qwen3.5-4B. For each problem we draw four samples and score the majority answer against the canonical output.

Method	Base model	Code Execution	Test Output	Avg.
Direct	Qwen3.5-2B	0.22 ± 0.041	0.19 ± 0.042	0.205 ± 0.042
Plan-and-Solve	Qwen3.5-2B	0.16 ± 0.039	0.10 ± 0.033	0.130 ± 0.031
Fixed-Strategy Ensemble	Qwen3.5-2B	0.19 ± 0.042	0.13 ± 0.034	0.160 ± 0.033
PKPO	Qwen3.5-2B	0.38 ± 0.050	0.21 ± 0.044	0.295 ± 0.042
CPPO	Qwen3.5-2B	0.41 ± 0.049	0.31 ± 0.047	0.360 ± 0.046
Direct	Qwen3.5-4B	0.38 ± 0.052	0.29 ± 0.044	0.335 ± 0.045
Plan-and-Solve	Qwen3.5-4B	0.15 ± 0.034	0.12 ± 0.030	0.135 ± 0.035
Fixed-Strategy Ensemble	Qwen3.5-4B	0.32 ± 0.048	0.30 ± 0.049	0.310 ± 0.042
PKPO	Qwen3.5-4B	0.48 ± 0.048	0.34 ± 0.048	0.410 ± 0.052
CPPO	Qwen3.5-4B	0.51 ± 0.054	0.35 ± 0.050	0.430 ± 0.052

International Conference on Learning Representations (ICLR).

Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fanjia Yan, Tianjun Zhang, Sida Wang, Armando Solar-Lezama, Koushik Sen, and Ion Stoica. 2025. LiveCodeBench: Holistic and contamination free evaluation of large language models for code. In *International Conference on Learning Representations (ICLR)*.

Tianjian Li, Yiming Zhang, Ping Yu, Swarnadeep Saha, Daniel Khashabi, Jason Weston, Jack Lanchantin, and Tianlu Wang. 2025. Jointly reinforcing diversity and quality in language model generations. *arXiv preprint arXiv:2509.02534*.

Yujia Li, David Choi, Junyoung Chung, Nate Kushman, Julian Schrittwieser, Remi Leblond, Tom Eccles, James Keeling, Felix Gimeno, Agustin Dal Lago, Thomas Hubert, Peter Choy, Cyprien de Masson d’Autume, Igor Babuschkin, Xinyun Chen, Po-Sen Huang, Johannes Welbl, Sven Gowal, Alexey Cherepanov, and 7 others. 2022. Competition-level code generation with AlphaCode. *Science*, 378(6624):1092–1097.

Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. 2024. Let’s verify step by step. In *International Conference on Learning Representations (ICLR)*.

Ilya Loshchilov and Frank Hutter. 2019. Decoupled weight decay regularization. In *International Conference on Learning Representations (ICLR)*.

Qwen Team. 2026. Qwen3.5 open models. Hugging Face model collection, <https://huggingface.co/Qwen>. Qwen3.5 series: 2B, 4B, and 9B base and chat variants. Exact checkpoint identifiers used in this paper are listed in Appendix E.6. Accessed 2026-05-26.

John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*.

Devan Shah, Owen Yang, Daniel Yang, Chongyi Zheng, and Benjamin Eysenbach. 2026. UpSkill: Mutual information skill learning for structured response diversity in LLMs. *arXiv preprint arXiv:2602.22296*.

Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. 2024. DeepSeekMath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*.

Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. 2025. Scaling LLM test-time compute optimally can be more effective than scaling model parameters. In *International Conference on Learning Representations (ICLR)*.

Christian Walder and Deep Karkhanis. 2025. Pass@k policy optimization: Solving harder reinforcement learning problems. *arXiv preprint arXiv:2505.15201*.

Evan Wang, Federico Cassano, Catherine Wu, Yunfeng Bai, Will Song, Vaskar Nath, Ziwen Han, Sean Hendryx, Summer Yue, and Hugh Zhang. 2024. Planning in natural language improves LLM search for code generation. *arXiv preprint arXiv:2409.03733*.

Lei Wang, Wanyu Xu, Yihuai Lan, Zhiqiang Hu, Yunshi Lan, Roy Ka-Wei Lee, and Ee-Peng Lim. 2023a. Plan-and-solve prompting: Improving zero-shot chain-of-thought reasoning by large language models. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2609–2634.

Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2023b. Self-consistency improves chain of thought reasoning in language models. In *International Conference on Learning Representations (ICLR)*.

Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L. Griffiths, Yuan Cao, and Karthik Narasimhan. 2023. Tree of thoughts: Deliberate problem solving with large language models. In *Advances in Neural Information Processing Systems (NeurIPS)*.

Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. 2023. Judging LLM-as-a-Judge with MT-Bench and chatbot arena. In *Advances in Neural Information Processing Systems (NeurIPS)*.

Yaoming Zhu, Sidi Lu, Lei Zheng, Jiaxian Guo, Weinan Zhang, Jun Wang, and Yong Yu. 2018. Texygen: A benchmarking platform for text generation models. In *Proceedings of the 41st International ACM SIGIR Conference on Research & Development in Information Retrieval*, pages 1097–1100.

A Main Results: Full Table Across All Base-Model Sizes

Table 8 reproduces Table 1 of the main text with the Qwen3.5-2B rows included. The 2B numbers track the 4B/9B trends already discussed in §1 but with smaller absolute deltas; the three unmarked cells—2B APPS, 4B LCBv6, and 9B APPS—are where the CPPO advantage stays positive but does not pass the hierarchical problem-and-seed bootstrap threshold.

Tuple Planner SFT samples the K plans jointly from $q_\Theta(s_i \mid x, s_{<i})$ (the joint-tuple policy class of §3.1, evaluated without RL); the matching iid-inference variant on the same checkpoint is reported in Appendix I as an ablation on autoregressive conditioning.

B Per-Seed Significance

The significance markers in Table 1 use a hierarchical bootstrap that resamples problems and training seeds (§4.3). Because three seeds give that test little power, we also report it directly: for each cell, Table 9 lists the CPPO-minus-strongest-baseline pass@4 difference computed separately on each of the three training seeds, their mean, and a 95% confidence interval over seeds (Student- t , two degrees of freedom). A cell is seed-significant when this interval excludes zero. The six cells whose interval excludes zero coincide with the six marked † in Table 1; the three borderline cells (2B APPS, 4B LCBv6, 9B APPS), where CPPO’s mean advantage is small and the strongest baseline is close, have intervals that include zero and are unmarked. The intervals are wide—an honest consequence of three seeds—but CPPO’s larger-margin wins survive training-seed resampling.

C Example Strategy Tuples from the Trained Planner

To make the trained planner’s behavior concrete, we show four strategy tuples emitted by the joint planner–solver policy (full CPPO) on held-out problems, together with the per-branch verifier outcome. Each box is one verbatim PLAN output: branch i is the solver attempt conditioned on method i , and ✓ marks a branch whose solution passed the verifier (unmarked branches failed). The tuples illustrate the planner proposing genuinely distinct algorithmic strategies for the same problem—brute-force enumeration, prefix-sum or

modular counting, and closed-form derivation—so that pass@ K succeeds whenever at least one strategy is both correct and solvable by the solver. Both problems are held-out CodeContests evaluation items (Table 3); the tuples are shown after final evaluation, purely for illustration, and are not used for model selection, tuning, or training.

Example 1 (LiveCodeBench abc367_d)

A: ✓ Compute prefix sums of A to determine clockwise distances, then iterate all pairs (s, t) while checking if the distance modulo M equals zero.
 B: Calculate the total sum of A modulo M and use modular arithmetic to derive the condition that the difference between prefix sums must be a multiple of M .
 C: Precompute prefix sums modulo M , then count pairs by iterating through the array and checking if the difference between current and previous prefix sums is zero.
 D: Use a frequency array to store counts of each prefix sum modulo M , then calculate the number of valid pairs by summing products of frequencies for matching remainders.

Example 2 (LiveCodeBench abc323_b)

A: ✓ Construct an adjacency matrix from the input strings to represent the tournament graph, then compute the win count for each player by summing the 0 entries in their corresponding row.
 B: ✓ Sort the players based on a custom key that prioritizes higher win counts first and breaks ties by ascending player index to ensure deterministic ranking.
 C: Implement a direct simulation where each player’s rank is calculated by comparing their win total against every other player’s win total and applying the tie-breaking rule explicitly.
 D: ✓ Use a frequency array to store win counts for $O(1)$ lookup during the sorting phase, avoiding repeated string traversals to optimize performance for larger N .

Table 8: Pass@4 results under the same $K=4$ solver-attempt budget and verifier across base model sizes. Full version of Table 1 of the main text, with the Qwen3.5-2B rows included. *Tuple Planner SFT* samples the K plans jointly from $q_\Theta(s_i \mid x, s_{<i})$; CPPO trains the joint planner-solver policy on this tuple class and gates planner credit on a learned plan-validity reward.

Method	Base model	APPS (Hendrycks et al., 2021)	CodeContests (Li et al., 2022)	LiveCodeBench (LCBv6) (Jain et al., 2025)
Direct Solve	Qwen3.5-2B	0.420 $\pm$ 0.017	0.040 $\pm$ 0.012	0.161 $\pm$ 0.014
Direct Solve	Qwen3.5-4B	0.515 $\pm$ 0.022	0.094 $\pm$ 0.014	0.214 $\pm$ 0.013
Direct Solve	Qwen3.5-9B	0.696 $\pm$ 0.021	0.175 $\pm$ 0.019	0.481 $\pm$ 0.014
Plan-and-Solve	Qwen3.5-2B	0.492 $\pm$ 0.021	0.065 $\pm$ 0.012	0.205 $\pm$ 0.014
Plan-and-Solve	Qwen3.5-4B	0.530 $\pm$ 0.019	0.098 $\pm$ 0.016	0.255 $\pm$ 0.012
Plan-and-Solve	Qwen3.5-9B	0.801 $\pm$ 0.015	0.171 $\pm$ 0.017	0.490 $\pm$ 0.020
PlanSearch	Qwen3.5-2B	0.513 $\pm$ 0.016	0.043 $\pm$ 0.012	0.225 $\pm$ 0.015
PlanSearch	Qwen3.5-4B	0.554 $\pm$ 0.022	0.128 $\pm$ 0.021	0.236 $\pm$ 0.012
PlanSearch	Qwen3.5-9B	0.770 $\pm$ 0.019	0.168 $\pm$ 0.018	0.564 $\pm$ 0.019
Pass@ K Training / RLVR	Qwen3.5-2B	0.476 $\pm$ 0.021	0.085 $\pm$ 0.015	0.215 $\pm$ 0.015
Pass@ K Training / RLVR	Qwen3.5-4B	0.607 $\pm$ 0.017	0.115 $\pm$ 0.020	0.316 $\pm$ 0.019
Pass@ K Training / RLVR	Qwen3.5-9B	0.762 $\pm$ 0.015	0.185 $\pm$ 0.021	0.503 $\pm$ 0.020
PKPO	Qwen3.5-2B	0.448 $\pm$ 0.024	0.062 $\pm$ 0.012	0.347 $\pm$ 0.016
PKPO	Qwen3.5-4B	0.722 $\pm$ 0.021	0.156 $\pm$ 0.020	0.488 $\pm$ 0.015
PKPO	Qwen3.5-9B	0.787 $\pm$ 0.020	0.183 $\pm$ 0.025	0.588 $\pm$ 0.015
UpSkill	Qwen3.5-2B	0.462 $\pm$ 0.018	0.041 $\pm$ 0.012	0.287 $\pm$ 0.018
UpSkill	Qwen3.5-4B	0.661 $\pm$ 0.023	0.121 $\pm$ 0.019	0.411 $\pm$ 0.016
UpSkill	Qwen3.5-9B	0.762 $\pm$ 0.015	0.136 $\pm$ 0.017	0.542 $\pm$ 0.013
Tuple Planner SFT	Qwen3.5-2B	0.453 $\pm$ 0.020	0.087 $\pm$ 0.020	0.232 $\pm$ 0.017
Tuple Planner SFT	Qwen3.5-4B	0.548 $\pm$ 0.024	0.147 $\pm$ 0.021	0.348 $\pm$ 0.018
Tuple Planner SFT	Qwen3.5-9B	0.733 $\pm$ 0.016	0.279 $\pm$ 0.026	0.514 $\pm$ 0.015
CPPO	Qwen3.5-2B	0.536 $\pm$ 0.017	0.201 $\pm$ 0.021	0.418 $\pm$ 0.017
CPPO	Qwen3.5-4B	0.784 $\pm$ 0.013	0.231 $\pm$ 0.025	0.505 $\pm$ 0.017
CPPO	Qwen3.5-9B	0.806 $\pm$ 0.019	0.396 $\pm$ 0.030	0.728 $\pm$ 0.016

[†] CPPO significantly outperforms the strongest non-CPPO baseline in that cell at $p < 0.05$ under a hierarchical bootstrap over problems and training seeds (1000 resamples; per-seed differences and 95% CIs in Appendix B). Unmarked cells—2B APPS, 4B LCBv6, 9B APPS—fall below the threshold; the CPPO advantage there is positive but the strongest baseline is already close (PlanSearch on 2B APPS, PKPO on 4B LCBv6, Plan-and-Solve at the 9B-APPS ceiling). *Baselines*. Baselines cover independent sampling, prompt-based planning, inference-time plan search, pass@ K -oriented RL, diversity-oriented RL, and planner-only SFT; implementation details are in §4.3 and Appendix R.

Table 9: Per-seed CPPO-minus-strongest-baseline pass@4 difference, with the mean and a 95% confidence interval over the three training seeds. Intervals excluding zero ($\dagger$) coincide with the significance markers in Table 1.

Base	Benchmark	δ_1	δ_2	δ_3	Mean	
Qwen3.5-2B	APPS	+0.010	+0.021	+0.038	+0.023	[−0.001, +0.046]
Qwen3.5-2B	CC	+0.100	+0.111	+0.131	+0.114 $\dagger$	[+0.001, +0.227]
Qwen3.5-2B	LCBv6	+0.058	+0.068	+0.087	+0.071 $\dagger$	[+0.001, +0.141]
Qwen3.5-4B	APPS	+0.049	+0.060	+0.077	+0.062 $\dagger$	[+0.001, +0.123]
Qwen3.5-4B	CC	+0.062	+0.072	+0.091	+0.075 $\dagger$	[+0.001, +0.149]
Qwen3.5-4B	LCBv6	+0.004	+0.015	+0.032	+0.017	[−0.001, +0.035]
Qwen3.5-9B	APPS	−0.008	+0.003	+0.020	+0.005	[−0.030, +0.040]
Qwen3.5-9B	CC	+0.194	+0.207	+0.232	+0.211 $\dagger$	[+0.100, +0.322]
Qwen3.5-9B	LCBv6	+0.127	+0.137	+0.156	+0.140 $\dagger$	[+0.050, +0.230]

Example 3 (CodeContests valid_66)

A: ✓ Generate the first 2000 integers sequentially, filter out those divisible by 3 or ending in 3, and store the valid numbers in a list to answer queries via direct indexing.

B: ✓ Precompute the sequence up to $k = 1000$ using a simple loop that increments a counter and checks divisibility by 3 or the last digit, storing results in an array for $O(1)$ retrieval.

C: Implement a direct formula approach by calculating the count of valid numbers

up to N as $\lfloor N/10 \rfloor + \lfloor (N+7)/10 \rfloor$ adjusted for the specific exclusion pattern, then binary search for the smallest N where this count equals k .

D: ✓ Use a mathematical transformation where each valid number maps to a 1-based index by skipping exactly one number for every three consecutive integers, then solve for the inverse mapping using modular arithmetic.

Example 4 (CodeContests valid_24)

A: Identify that x is interesting if and only if x ends in one or more 9s, then count how many such numbers exist up to n by iterating through the number of trailing 9s and using digit DP or combinatorics to count valid prefixes.

B: Implement a digit DP where the state tracks the current position, whether we are already below the limit n , and the count of trailing 9s seen so far, summing up valid states where the final number has at least one trailing 9.

C: Observe that the number of interesting integers up to n is exactly the sum of $\lfloor n/10^k \rfloor$ for $k \geq 1$, which can be com-

puted in $O(\log n)$ time by iterating powers of 10 without simulating each number.

D: ✓ Use a direct mathematical formula derived from the property that every interesting number is of the form $d \dots d99 \dots 9$, counting occurrences by summing the contribution of each power of 10 to the total count up to n .

D Planner Prompts

Three prompts govern the planner and solver, shown below: the PLAN contract and the SOLVE prompt used at training and inference, and the candidate-generation prompt used to sample the planner SFT data. The PLAN and candidate-generation prompts impose a strict conciseness requirement on the generated strategies; in the token-accounting slices of Table 4, this concise strategy condition is associated with shorter solver completions, offsetting the additional planner generation in decoded-token accounting.

PLAN prompt contract (training and inference)

The planner emits a strategy tuple of K alternative methods. Methods are labelled A:, B:, C:, D: on separate lines and must obey the following contract.

- Each attempt must be exactly one concise sentence on its own label line.
- Keep each attempt short: no more than 45 words per label.
- Do not include substeps, derivations, implementation walkthroughs, or multi-paragraph methods.
- Stop immediately after the D: method; do not append any extra analysis or commentary.

SOLVE prompt (training and inference)

The solver receives the problem statement and a single strategy s_i from the plan tuple, and returns one complete solution that follows that strategy.

- Implement the given strategy faithfully; do not switch to a different approach or merge in other methods.
- Output a single self-contained program in the target language, enclosed in one code block.

- Emit only the program—no explanation, no alternative solutions, and no commentary outside the code block.
- The program is scored solely by the sandboxed verifier against the official tests (§4.3).

SFT candidate-generation prompt (base planner, 6 candidates per problem)

Generate 6 candidate high-level solver methods for this programming problem.

Candidate method rules.

- Each method must be independently usable by a solver.
- Each method must be specific to this exact problem.
- Each method must include a concrete hook—the data structure, invariant, or transformation that distinguishes it from the others.
- Prefer genuinely different strategies or useful variants over cosmetic rewrites of the same approach.
- Be concise and precise. Do not include filler, generic advice, wrong-task methods, unsupported assumptions, cosmetic rewrites, code, pseudocode, final answers, or full solution traces.

For each training problem we sample candidate methods from the base planner with the prompt above; a strict Qwen3.5-9B judge then admits only candidates that satisfy the above rules; only problems for which at least K strict-clean candidates survive enter the SFT set, after which we generate the SFT trace from the selected K methods. Stage-by-stage candidate counts appear in Appendix E.4.

E Data Splits and Decontamination

This appendix gives the full per-component exposure audit, the decontamination algorithm and cross-corpus overlap counts, the candidate selection funnel, and the sandbox configuration referenced from §4.2 and §4.3.

E.1 Stage-wise splits with problem counts

Table 10 reports the problem counts behind Table 3 of the main text. The top block covers every training-side pool: planner SFT trains on strict $K=4$ Think+Plan gold tuples filtered from Code-Contests train; the reward model is trained and

validated on a balanced Pass/No-Pass subset of the same corpus; the planner warm-up also uses CodeContests train to learn the plan contract before joint training; and the joint CPPO rollout pool draws from CodeContests train. The bottom block lists the held-out evaluation sets reported in Table 1; each is disjoint from every training-side row under the decontamination protocol of Appendix E.3.

Table 10: Stage-wise data splits with problem counts. Reward-model counts are reported as “train + val”, balanced 1:1 Pass/No-Pass. The bottom block lists the held-out evaluation sets reported in Table 1; each is disjoint from every training-side row in the top block.

Stage	Source split	# Problems
Planner (Think+Plan gold)	SFT CC train	333
Reward-model train + val	CC train (filtered)	1110 + 278
Planner warm-up	CC train (filtered)	500
Joint CPPO rollout pool	CC train	300
Held-out eval — APPS	APPS test (intro)	200
Held-out eval — CodeContests	CC valid (official)	100
Held-out eval — LCBv6	LCBv6 held-out	300

E.2 Exposure audit

Table 11 lists every component that consumes problem statements, with the source split it queries and whether any held-out evaluation problem was ever observed. LCBv6 held-out evaluation prompts appear in no row.

Evaluation slices. When a diagnostic uses a subset rather than a full benchmark split, the slice is fixed before model comparison and shared by all methods. The decoded-token accounting in Appendix M uses the held-out APPS and LiveCodeBench-v6 evaluation sets of Table 10. Problem IDs for all final evaluation and diagnostic slices are released with the run logs; no slice is chosen based on model outputs.

E.3 Decontamination algorithm

Deduplication uses normalized prompt-hash matching: we lowercase the statement, strip whitespace and HTML tags, normalize $\text{\LaTeX}$ (math symbols, environments) to a canonical form, remove sample I/O blocks, and hash the result with SHA-256. Source-level identifiers (Codeforces contest + problem index, APPS id, LCBv6 id) are used when both items expose them. Under this protocol, every cross-corpus train $\leftrightarrow$ eval pair (APPS $\leftrightarrow$ CC, APPS $\leftrightarrow$ LCBv6, CC $\leftrightarrow$ LCBv6) has zero prompt-hash overlap. CodeContests train and held-out additionally share split-local synthetic identifiers from

our id assignment; those identifier collisions do not correspond to repeated problem content under prompt-hash.

E.4 Candidate selection funnel

We construct the planner SFT set with an over-generate-then-filter pipeline. For each training problem, the base planner samples multiple candidate plans; a Qwen3.5-9B judge keeps only methods that are problem-specific, solver-actionable, code-free, and non-duplicate. A problem enters the SFT set only if at least K such methods remain, after which we generate the final SFT trace from the selected K methods. This filter selects for parseable, multi-strategy problems rather than frozen-solver success, so it biases the SFT data toward problems with several admissible approaches, not toward problems that are easy for the solver. The resulting planner SFT set uses only CodeContests-train problems; no held-out evaluation problem enters the sampling or judge stages (Table 11).

E.5 Sandbox configuration

Every candidate program is run in an isolated sandbox with the limits in Table 12. Compilation error, runtime error, timeout, memory cap, and empty extraction are all scored as $V(x, y) = 0$ at the verifier; the same scoring rule applies to training rollouts and final evaluation.

E.6 Model checkpoints

Table 13 lists the exact model identifier and the base/instruction-tuned designation for every checkpoint used in the paper. The planner-solver backbones and the plan-validity reward model are *base* checkpoints—CPPO performs its own SFT and RL on top of them—while the offline judge and the Gemma 4 generality models are *instruction-tuned*. All weights are loaded from the official Hugging Face releases of Qwen Team (2026) and Google DeepMind (2026); the pinned revision hash for each repository is recorded in the released configuration files (rollout recipe, Appendix Q).

F Reward-Model Training Details

This appendix gives the exact splits, validation thresholds, and refresh schedule for the plan-validity reward model referenced in Stage 2 of §3.3.

Labeled-tuple pool and splits. The CodeContests-only RM package contains

Table 11: Per-component exposure audit. “Held-out eval seen?” covers every held-out evaluation split listed in Table 3.

Component	Source split(s)	Held-out eval seen?	Notes
Plan candidates (base policy)	CC train	no	self-generated
LLM judge (Qwen3.5-9B) for plan validity	CC train	no	top- K selection only
RM training labels	CC train (filtered)	no	balanced pass/fail
RM val (early stopping, threshold)	CC train (filtered)	no	disjoint from RM train
RM refresh during CPPO	CPPO rollout pool	no	same source as CPPO
Planner warm-up	CC train	no	500 problems
Joint CPPO rollout	CC train	no	no held-out eval problems
Hyperparameter sweep / early stopping	CC train (dev subset)	no	dev disjoint from eval
Qualitative examples (Appendices C, G)	CC valid (held-out eval)	N/A	post-evaluation illustration; no updates or selection
Final evaluation (Table 1)	APPS / CC / LCBv6 held-out	N/A	no parameter updates

Table 12: Sandbox configuration used by the verifier (training rollouts and final evaluation share the same limits).

Resource	Limit
Wall-clock timeout (per test case)	10 s
Memory	256 MiB per subprocess (RLIMIT_AS, fallback RLIMIT_RSS)
Process cap	RLIMIT_NPROC = 64
Network access	disabled
Filesystem	read-only image, ephemeral scratch
Available standard library	Python 3 + numpy, no third-party packages

Table 13: Exact model checkpoints. All weights come from the official Hugging Face releases; per-repository revision hashes are pinned in the released configs.

Role	Hugging Face repository	Variant
Planner–solver backbone (2B)	Qwen/Qwen3.5-2B-Base	base
Planner–solver backbone (4B)	Qwen/Qwen3.5-4B-Base	base
Planner–solver backbone (9B)	Qwen/Qwen3.5-9B-Base	base
Plan-validity reward model	Qwen/Qwen3.5-0.8B-Base	base
Offline judge (SFT filter, diversity)	Qwen/Qwen3.5-9B-Instruct	instruct
Generality check	google/gemma-4-e2b-it	instruct
Generality check	google/gemma-4-e4b-it	instruct

1388 balanced tuple labels: 1110 train labels (555 Pass / 555 No-Pass) and 278 held-out validation labels (139 Pass / 139 No-Pass). These labels are built from strict-four judge-validated plans and semantically invalid but superficially well-formed negatives, then downsampled to a 1:1 Pass/No-Pass ratio. All reward-model labels come from CodeContests train; no LiveCodeBench problem is ever judged or labeled.

Pre-registered validation thresholds. We accept a reward-model checkpoint ψ only when its held-out validation metrics jointly satisfy $\text{AUC} \geq 0.75$, balanced accuracy ≥ 0.70 , and precision and recall both ≥ 0.65 at the decision threshold of §4.1. We exclude raw accuracy from the acceptance criterion because, at a single fixed threshold, it does not reflect performance at the high-recall operating point ($\tau=0.17$, §4.1) at which the gate is actually used; AUC, balanced accuracy, and threshold-specific precision and recall capture this better.

Refresh schedule. After each CPPO phase we relabel a fresh batch of planner outputs with the same judge, rebalance by Pass / Fail and by prompt, and resume fine-tuning ψ from its latest checkpoint. Freezing the gate within each phase keeps optimization stable; refreshing between phases keeps the gate aligned with the shifted planner distribution.

G Qualitative Example: Strategy Tuples

Table 14 contrasts the $K=4$ attempts from CPPO with those from independent base-model sampling (Direct Solve) at the same budget on a representative problem. All eight attempts compile and run; the difference is algorithmic. CPPO’s planner emits four distinct strategies, and the correct one (two DSUs, one per forest) passes the verifier. Direct Solve’s four iid samples all compile cleanly but converge on the same incorrect single-DSU greedy—the dominant mode of the base model’s answer distribution on this problem—so every branch fails the same way and pass@4 collapses to 0.

H Plan-Diversity Measurement

We measure pairwise method diversity to distinguish coordinated strategy coverage from a planner that merely paraphrases one approach. We make that signal precise below, validate its most consequential form against human labels, and report per-method values in Table 16—turning the strategy-coverage claim of §1 from illustration into measurement.

Table 14: Illustrative $K=4$ attempts on Codeforces 1559D1 (Mocha and Diana, Easy Version), drawn from the held-out CodeContests evaluation split. CPPO’s left-column rows are planner-emitted strategy sketches; Direct Solve’s right-column rows summarize the implemented algorithm of each iid solver sample (all four compile and execute against the official tests). The verifier outcome ($\checkmark$ / $\times$) is the per-branch sandboxed test result.

	CPPO (coordinated tuple)		Direct Solve (iid base-model samples)
Problem	Two players each hold a forest on the same $n \leq 1000$ vertices. Add the same set of edges to both graphs so that both remain forests. Output the maximum number of added edges and any valid edge set. Reference solution maintains one DSU per forest and greedily adds an edge (u, v) iff u and v are disconnected in both DSUs.		
1	Pairwise DSU scan over all (u, v) pairs.	$\times$	Single-DSU greedy, iterative path compression; scans (u, v) pairs and adds (u, v) iff disconnected in the merged DSU. $\times$
2	Greedy edge addition with two DSUs: add (u, v) iff u, v are disconnected in both DSUs, then union in both. $\checkmark$	$\times$	Single-DSU greedy with union-by-rank; identical edge-add rule, different union heuristic. $\times$
3	Incremental DSU construction with a candidate-priority can_add predicate. $\times$	$\times$	Single-DSU greedy with recursive find; otherwise the same algorithm. $\times$
4	Graph-matching view: component abstraction matched across the two forests. $\times$	$\times$	Single-DSU greedy with adjacency-list bookkeeping; same merge rule. $\times$
pass@4	$\checkmark$ (branch s_2 passes)	$\times$ (all four implement the same incorrect single-DSU greedy and fail on the same dual-forest counterexample)	

H.1 Definitions

For a K -method plan tuple $S = (s_1, \dots, s_K)$ produced for prompt x , we measure diversity at three levels: surface, semantic, and algorithmic. Methods without an explicit plan—Direct Solve, Pass@ K -only RL, and PKPO—have no S ; for those we apply the same three metrics to the K solver answers $(y_1, \dots, y_K)$ instead, so every method row of Table 1 has a comparable diversity number.

Surface diversity (D_{surf}). We use one minus mean pairwise Self-BLEU (Zhu et al., 2018) on plan tokens,

$$D_{\text{surf}}(S) = 1 - \frac{1}{K(K-1)} \sum_{i \neq j} \text{BLEU}_4(s_i, s_j), \quad (14)$$

so that $D_{\text{surf}} = 1$ when the K method texts are pairwise disjoint at the 4-gram level and 0 when they are token-identical. Surface diversity catches verbatim and near-verbatim duplicates but is fooled by superficial rewrites of one underlying algorithm (e.g. “top-down DP” vs. “memoized recursion”).

Semantic diversity (D_{sem}). We encode each s_i into $\mathbb{R}^{768}$ with the all-mpnet-base-v2 sentence encoder $e(\cdot)$. Writing $d_{\cos}(u, v) = 1 - \cos(u, v)$ for the cosine distance, we take the mean pairwise

distance over plan embeddings,

$$D_{\text{sem}}(S) = \frac{1}{K(K-1)} \sum_{i \neq j} d_{\cos}(e(s_i), e(s_j)). \quad (15)$$

Semantic diversity is harder to fool with paraphrases than D_{surf} , but a small D_{sem} gap can still underweight algorithmic distinctions when two short noun-phrase methods embed close (“DP” and “greedy”).

Algorithmic diversity (D_{alg}). We fix a taxonomy $\mathcal{T}$ of common algorithmic strategies in competitive programming—DP, greedy, graph, search/backtrack, number-theory/math, simulation, data-structure, divide-and-conquer, brute-force, and other—and map each method s_i to a category $c_i \in \mathcal{T}$ by prompting Qwen3.5-9B as an offline LLM-as-judge (Zheng et al., 2023) with a zero-shot template that lists the taxonomy and asks for the single best-matching label. Algorithmic diversity is the unique-category coverage,

$$D_{\text{alg}}(S) = \frac{|\{c_1, \dots, c_K\}|}{K} \in (0, 1], \quad (16)$$

so $D_{\text{alg}} = 1$ exactly when the K methods land in K distinct categories of $\mathcal{T}$, and $D_{\text{alg}} = 1/K$ when they collapse onto one. This is the metric most directly aligned with the strategy-coverage claim of §1: it cannot be inflated by rewording, and it is invariant to which surface tokens or which

embedding model is used. For methods without an explicit plan, the same classifier is applied directly to the solver code with a prompt asking which algorithm class the program implements.

H.2 Reliability of the algorithmic-category classifier

Because D_{alg} is the metric most directly tied to the strategy-coverage claim, we validate the classifier against human labels on a small balanced held-out set. Two annotators independently label each plan against $\mathcal{T}$ from the prompt and the plan text alone. We report inter-annotator agreement (Cohen’s κ) and the classifier’s accuracy and macro-F1 against the consensus label in Table 15. We treat D_{alg} as informative only if Cohen’s $\kappa \geq 0.6$ and macro-F1 ≥ 0.7 ; failing either threshold, the classifier is replaced or the taxonomy is collapsed before any cell of Table 16 is reported.

Table 15: Reliability of the algorithmic-category classifier of (16), evaluated on a balanced held-out set against two-annotator consensus labels over the taxonomy $\mathcal{T}$.

Quantity	Value
Cohen’s κ (inter-annotator)	0.72
Classifier accuracy	0.81
Classifier macro-F1	0.78

H.3 Diversity across methods

Table 16 reports the three diversity metrics at $K=4$ on APPS held-out problems with Qwen3.5-2B as the base model, under the same evaluation setting as the headline APPS result. For each method, diversity is averaged first over problems and then over three training seeds, with the same seed protocol as §4.3. A coordinated-strategy account predicts a clean separation: D_{alg} should rank CPPO above every iid-sampling baseline (Direct Solve, Pass@ K -only RL, PKPO) regardless of pass@ K , since each of those baselines draws its K samples from a single answer distribution.

H.4 Diversity vs. pass@ K gain

Figure 4 relates per-cell algorithmic diversity D_{alg} at $K=4$ to the absolute pass@4 improvement each method achieves over Direct Solve on the same (benchmark, base model) pair. The positive slope supports an association between algorithmic diversity and pass@ K gain: methods whose $K=4$ attempts span more algorithmic strategies tend to

Table 16: Plan-tuple diversity at $K=4$ on APPS / Qwen3.5-2B, mean $\pm$ std over three training seeds. D_{surf} uses Self-BLEU, D_{sem} uses sentence embeddings, and D_{alg} uses the algorithmic-category classifier of (16). For methods without an explicit plan, we compute the analogous diversity over the generated solver answers.

Method	D_{surf}	D_{sem}	D_{alg}
Direct Solve	0.74 $\pm$ 0.04	0.31 $\pm$ 0.03	0.38 $\pm$ 0.05
Plan-and-Solve	0.77 $\pm$ 0.04	0.34 $\pm$ 0.04	0.44 $\pm$ 0.06
PlanSearch	0.83 $\pm$ 0.03	0.41 $\pm$ 0.04	0.58 $\pm$ 0.05
Pass@ K Training / RLVR	0.76 $\pm$ 0.05	0.33 $\pm$ 0.04	0.43 $\pm$ 0.06
PKPO	0.79 $\pm$ 0.04	0.37 $\pm$ 0.04	0.49 $\pm$ 0.06
Tuple Planner SFT	0.86 $\pm$ 0.03	0.47 $\pm$ 0.04	0.66 $\pm$ 0.05
CPPO	0.90$\pm$0.03	0.55$\pm$0.04	0.79$\pm$0.04

improve more over independent sampling. Diversity is necessary but not sufficient, however—a handful of cells combine high D_{alg} with near-zero gain, the four-distinct-but-uniformly-wrong regime. Diversity-enhancing methods (PKPO, UpSkill, CPPO) cluster in the upper-right of the panel, while prompting-only baselines (Plan-and-Solve, fixed-skill prompting) sit in the lower-left. We position this figure as a diagnostic of the relationship rather than a causal proof: the causal claim is made by the planner ablations of §4.5.1, which remove the diversity-inducing components one at a time.

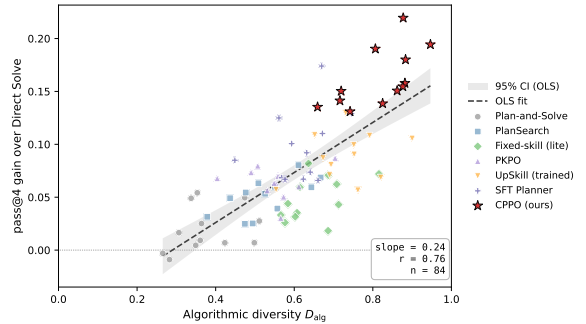


Figure 4: Relationship between algorithmic diversity and pass@4 improvement. Each point is a (method, benchmark, base-model) cell, excluding Direct Solve which serves as the zero-gain reference. The x-axis shows D_{alg} at $K = 4$; for planner-based methods it is computed over the four planned strategies, and for no-planner baselines it is computed over the four generated solutions. The y-axis reports the absolute pass@4 gain over Direct Solve under the same benchmark and base model. The dashed line is an OLS fit with a 95% confidence band. Additional robustness checks with alternative diversity metrics are reported in Appendix H.

I Inference-Mode Ablation: Joint vs. iid Sampling

The joint policy-class claim in §3.1 rests on autoregressive conditioning: each plan s_i is sampled from $q_\Theta(s_i \mid x, s_{<i})$, so later strategies can avoid duplicating earlier ones. To separate the contribution of this inference-time mechanism from the training-time effect of the SFT and RL stages, we evaluate the same checkpoints under *iid inference*: the K plans are sampled independently from $q_\Theta(\cdot \mid x)$, clearing the planner’s history between draws. All other settings—prompt, sampling temperature, top- p , verifier, and $K=4$ budget—are held fixed; iid sampling uses the same hyperparameters as joint inference and is not tuned separately, and plans are drawn with independent random seeds to rule out degenerate identical draws.

Table 17 reports pass@4 for Tuple Planner SFT and Full CPPO under both inference modes at three model sizes; Joint columns reproduce the corresponding Tuple Planner SFT and CPPO rows of Table 1, and iid columns evaluate the same checkpoints with the planner’s history cleared between the K plan draws. CPPO consistently loses more pass@4 under iid inference than Tuple Planner SFT does (Δ of 0.032–0.053 for CPPO versus 0.012–0.018 for SFT), and the effect attenuates with model scale—consistent with the view that larger bases produce marginally more diverse plans, so autoregressive conditioning has less unique work to do at 9B than at 2B.

Table 17: Inference-mode ablation on APPS, $K=4$. *Joint* inference samples plans autoregressively from $q_\Theta(s_i \mid x, s_{<i})$; *iid* inference samples each plan independently from $q_\Theta(\cdot \mid x)$ by clearing the planner’s history between draws. Both modes share the same checkpoint, sampling temperature, and verifier budget; only the inference loop differs. Joint values reproduce Table 1.

Base model	Checkpoint	Joint	iid	Δ
Qwen3.5-2B	Tuple Planner SFT	0.453 $\pm$ 0.020	0.435 $\pm$ 0.023	+0.018
	Full CPPO	0.536 $\pm$ 0.017	0.483 $\pm$ 0.022	+0.053
Qwen3.5-4B	Tuple Planner SFT	0.548 $\pm$ 0.024	0.531 $\pm$ 0.021	+0.017
	Full CPPO	0.784 $\pm$ 0.013	0.731 $\pm$ 0.019	+0.053
Qwen3.5-9B	Tuple Planner SFT	0.733 $\pm$ 0.016	0.721 $\pm$ 0.018	+0.012
	Full CPPO	0.806 $\pm$ 0.019	0.774 $\pm$ 0.020	+0.032

Table 18 reports the matching algorithmic diversity D_{alg} (Appendix H, (16)) on Qwen3.5-2B / APPS. The diversity drop under iid inference is larger for CPPO (−0.19) than for Tuple Planner SFT (−0.12), so the pass@4 gap in Table 17 tracks

an actual reduction in strategy coverage rather than a sampling artifact. Together, the two tables show that autoregressive conditioning at inference time is a substantive and method-specific contributor to CPPO’s coordination gain: the RL stage trains the planner to use this conditioning, and removing it at inference partially undoes the gain.

Table 18: Algorithmic diversity D_{alg} at $K=4$ on Qwen3.5-2B / APPS under each inference mode. Joint values reproduce Table 16.

Checkpoint	Joint D_{alg}	iid D_{alg}	Δ
Tuple Planner SFT	0.66 $\pm$ 0.05	0.54 $\pm$ 0.06	+0.12
Full CPPO	0.79 $\pm$ 0.04	0.60 $\pm$ 0.05	+0.19

Decomposing CPPO’s pass@4 gain. Subtracting Direct Solve from the iid and joint columns of Table 17 separates two sources of gain that the joint policy class actually combines: *training-induced diversity*, the improvement over Direct Solve that survives under iid inference (CPPO iid – Direct Solve), and *inference-time autoregressive conditioning*, the further improvement from sampling jointly on the same checkpoint (CPPO joint – CPPO iid). Table 19 reports both at each model size. The training-induced share is the larger of the two across all three sizes (54–80% of the total CPPO gain), and the inference-time share is smaller and shrinks at 9B (+0.032, versus +0.053 at 2B and 4B). Two readings are consistent with this pattern: the joint policy class operates primarily through training—the CPPO loss is defined over $q_\Theta(s_i \mid x, s_{<i})$, so the planner learns to allocate probability mass across distinct strategies even when later evaluated iid—and the residual inference-time gain captures the additional benefit of keeping the conditioning at sampling time. The 9B contraction is consistent with a larger marginal $q_\Theta(\cdot \mid x)$ that already covers a wider strategy distribution, leaving autoregressive conditioning less unique work to do at scale; the training-induced share (+0.078 at 9B) remains substantial.

Table 19: Decomposing CPPO’s pass@4 gain over Direct Solve on APPS into a training-induced share (CPPO iid – Direct Solve) and an inference-time share (CPPO joint – CPPO iid). Direct Solve and CPPO values are reproduced from Table 8 and Table 17.

Base model	Total gain	Training share	Joint-inference share	Inference %
Qwen3.5-2B	+0.116	+0.063	+0.053	46%
Qwen3.5-4B	+0.269	+0.216	+0.053	20%
Qwen3.5-9B	+0.110	+0.078	+0.032	29%

Order-randomized control for truncation.

Stage 1 (§3.3) uses the $K_{\text{tuple}}=4$ planner throughout the paper, so the small- K rows of Table 6 take the first K_{solve} outputs of one 4-tuple. To check that this introduces no positional bias, we evaluate Full CPPO at $K_{\text{solve}}=2$ on Qwen3.5-2B / APPS under three selection rules: *first 2* (the rule used in the main text), *last 2*, and *uniformly random 2 of 4*. The three estimates lie within one seed-std of each other— 0.518 ± 0.048 , 0.514 ± 0.049 , 0.516 ± 0.048 —supporting position-independence in practice for this truncation rule.

J GRPO Optimizer Details

For completeness, we spell out the standard GRPO objective used by CPPO. For each prompt, GRPO samples G rollouts $\{y_1, \dots, y_G\}$ from $\pi_{\theta_{\text{old}}}$, scores them with rewards r_i , and computes group-normalized advantages

$$A_i = \frac{r_i - \text{mean}(r_{1:G})}{\text{std}(r_{1:G}) + \epsilon}. \quad (17)$$

The clipped, KL-regularized objective is

$$\begin{aligned} \mathcal{J}_{\text{GRPO}}(\theta) = \mathbb{E}[\min(\rho_i A_i, \text{clip}(\rho_i, 1 \pm \epsilon) A_i)] \\ - \beta \text{KL}(\pi_{\theta} \parallel \pi_{\text{ref}}), \end{aligned} \quad (18)$$

where $\rho_i = \pi_{\theta}(y_i \mid x) / \pi_{\theta_{\text{old}}}(y_i \mid x)$ is the importance ratio. CPPO uses this optimizer unchanged.

K Pseudocode for the Full CPPO Pipeline

Algorithm 2 writes out all five stages of the end-to-end pipeline of §3.3: planner SFT, reward-model training, RM-guided planner warm-up, the audit gate, and the joint CPPO loop whose inner step is the one summarized in §3.2.

L RM-Guided Planner Warm-Up

Why planner warm-up helps. The RM-guided warm-up is not intended to directly improve solver accuracy. During this stage the solver is frozen and no verifier outcome is used; the planner is optimized only against the validity gate,

$$R_{\text{warm}}(x, S) = J_{\psi}(x, S). \quad (19)$$

Its role is to align the planner’s on-policy distribution with the frozen plan-validity gate before joint CPPO. This alignment matters because the planner reward in CPPO is multiplicative,

$$R_{\text{plan}}(x, S) = J_{\psi}(x, S) \cdot R_{\text{out}}(x, S), \quad (20)$$

so when the planner frequently samples tuples rejected by J_{ψ} , solver success does not yield planner credit: $R_{\text{plan}} = 0$ even when $R_{\text{out}} = 1$. Warm-up reduces these invalid rollouts and raises the density of non-zero gated outcome rewards, which stabilizes the subsequent joint planner–solver update. Warm-up therefore provides only a small standalone gain under a frozen solver, but it gives the joint CPPO stage a denser reward signal.

M Decoded-Token Accounting

Table 20 gives the decoded-token accounting for the main-text token-normalized slices. Planner and solver tokens are separated where the run logs expose that split: planner tokens count generated plans or strategy tuples, and solver tokens count generated solutions. Prompt tokens are excluded here, matching Table 4. The planner–solver split is unavailable for the LCBv6 PlanSearch and CPPO runs (the harness recorded only aggregate decoded length); the corresponding cells are dashed, but the total decoded-token counts and the $\text{pass}@K/10k$ values are complete.

N LiveCodeBench-v6 Pass@K Sweep

Figure 5 extends Figure 2 of the intro with the Qwen3.5-4B panel alongside the Qwen3.5-9B panel already shown in the intro. Both panels share the same $K \in \{1, 2, 4, 8, 16\}$ sweep on LiveCodeBench-v6. CPPO follows an S-shape on both sizes: at $K=1$ the planner overhead leaves it at-or-below the independent-sampling baselines; coordination starts to pay off from $K \geq 2$, the curve overtakes PKPO (the strongest non-CPPO baseline at $K=4$ on both sizes) by $K=4$, and saturates from $K=8$ onward. The inflection sits earlier on the weaker 4B base ($K=1 \rightarrow 2$) than on the 9B ($K=2 \rightarrow 4$): a stronger solver makes a single strategy more often sufficient at small budgets, delaying the point at which a second alternative starts to pay off.

O Generality across Base Models: Gemma 4

The main results in Table 1 use the Qwen3.5 family throughout. To check that CPPO’s gains are not tied to a specific backbone, we rerun the full pipeline on Gemma 4 E2B and E4B (Google DeepMind, 2026) without retuning any component. The training data follows the Qwen3.5 protocol: CodeContests supplies planner SFT, reward-model

Algorithm 2 CPPO end-to-end training pipeline

Require: base policy Θ , prompts $\mathcal{D}$, offline judge $\mathcal{J}$, verifier V , K branches, M planner samples per prompt, step budgets $T_{\text{SFT}}, T_{\text{RM}}, T_{\text{wu}}, T_{\text{cpo}}$

Stage 1: Planner SFT ▷ §3.3

- 1: **for** $t = 1, \dots, T_{\text{SFT}}$ **do**
- 2: Train planner tokens of Θ by cross-entropy on offline self-generated, judge-validated strict-four gold tuples $\{(x, S^*)\}_{x \in \mathcal{D}_{\text{gold}}}$.
- 3: **end for**

Stage 2: Reward-model training ▷ §3.3

- 4: Sample (x, S) from Θ ; label with $\mathcal{J}$; balance by prompt and label.
- 5: **for** $t = 1, \dots, T_{\text{RM}}$ **do**
- 6: Update ψ by BCE: $\mathcal{L}_{\text{RM}} = -z \log \hat{z}_\psi - (1-z) \log(1-\hat{z}_\psi)$.
- 7: **end for**
- 8: Accept ψ only if held-out AUC, balanced accuracy, and precision/recall all pass.

Stage 3: RM-guided planner warm-up ▷ §3.3

- 9: **for** $t = 1, \dots, T_{\text{wu}}$ **do**
- 10: For each x , sample $S \sim q_\Theta(\cdot | x)$; form $R_{\text{warm}}(x, S) = J_\psi(x, S)$.
- 11: Apply clipped, KL-regularized GRPO (18) to planner tokens only.
- 12: **end for**

Stage 4: Audit (no parameter update)

- 13: Measure frozen-solver pass@ K and the fraction of rollouts with nonzero R_{plan} ; if either is insufficient, return to Stage 1 or 2.

Stage 5: Joint CPPO ▷ §3.2

- 14: **for** $t = 1, \dots, T_{\text{cpo}}$ **do**
 - 15: **for all** prompts x in the batch **do**
 - 16: Sample M planner tuples $S^{(m)} \sim q_\Theta(\cdot | x)$; score $J_\psi(x, S^{(m)})$.
 - 17: For each $s_i^{(m)}$, sample $y_i^{(m)} \sim p_\Theta(\cdot | x, s_i^{(m)})$; verify $r_i^{(m)} = V(x, y_i^{(m)})$.
 - 18: Within-tuple normalize $\{r_i^{(m)}\} \rightarrow$ solver advantages $a_i^{(m)}$ (11).
 - 19: $R_{\text{out}}^{(m)} \leftarrow \mathbb{I}\{\max_i r_i^{(m)} = 1\}$; $R_{\text{plan}}^{(m)} \leftarrow J_\psi(x, S^{(m)}) \cdot R_{\text{out}}^{(m)}$.
 - 20: Across-tuple normalize $\{R_{\text{plan}}^{(m)}\} \rightarrow$ planner advantages $A_{\text{plan}}^{(m)}$ (12).
 - 21: **end for**
 - 22: Apply GRPO solver loss with $a_i^{(m)}$ and GRPO planner loss with $A_{\text{plan}}^{(m)}$; update Θ via (13).
 - 23: Optionally refresh ψ on a fresh batch of planner outputs labeled by the same judge (Stage 2, §3.3).
 - 24: **end for**
 - 25: **return** Θ
-

Table 20: Decoded-token accounting for Qwen3.5-4B pass@4 slices. All values are averaged per problem except pass@ K and pass@ $K/10k$ tokens.

Model	Dataset	Method	K	Planner toks.	Solver toks.	Total decoded toks.	pass@ K	pass@ $K/10k$ toks.
Qwen3.5-4B	APPS	Direct Solve	4	0.0	6086.8	6086.8	0.515	0.846
Qwen3.5-4B	APPS	Plan-and-Solve	4	0.0	7326.2	7326.2	0.530	0.723
Qwen3.5-4B	APPS	PlanSearch	4	566.0	4969.0	5535.0	0.554	1.001
Qwen3.5-4B	APPS	CPPO	4	131.1	4596.7	4727.8	0.784	1.658
Qwen3.5-4B	LCBv6	Direct Solve	4	0.0	1731.0	1731.0	0.214	1.236
Qwen3.5-4B	LCBv6	Plan-and-Solve	4	0.0	3975.0	3975.0	0.255	0.642
Qwen3.5-4B	LCBv6	PlanSearch	4	–	–	2137.0	0.236	1.104
Qwen3.5-4B	LCBv6	PKPO	4	0.0	2137.0	2137.0	0.488	2.284
Qwen3.5-4B	LCBv6	UpSkill	4	0.0	2218.0	2218.0	0.411	1.853
Qwen3.5-4B	LCBv6	CPPO	4	–	–	1556.0	0.505	3.246

data, planner warm-up, and joint CPPO rollouts; APPS and LiveCodeBench-v6 remain evaluation-only. The planner/solver roles, RM warm-up, GRPO hyperparameters, $K=4$ budget, verifier, and evaluation protocol are otherwise identical to the Qwen3.5 runs in §4.3; we only swap in Gemma 4

base weights for both planner and solver. Table 21 reports pass@4 on APPS and LiveCodeBench-v6. CPPO achieves the highest pass@4 at both sizes on both benchmarks, matching the pattern from the Qwen3.5 sweep in Table 1. The Gemma 4 base is already much stronger on APPS than Qwen3.5

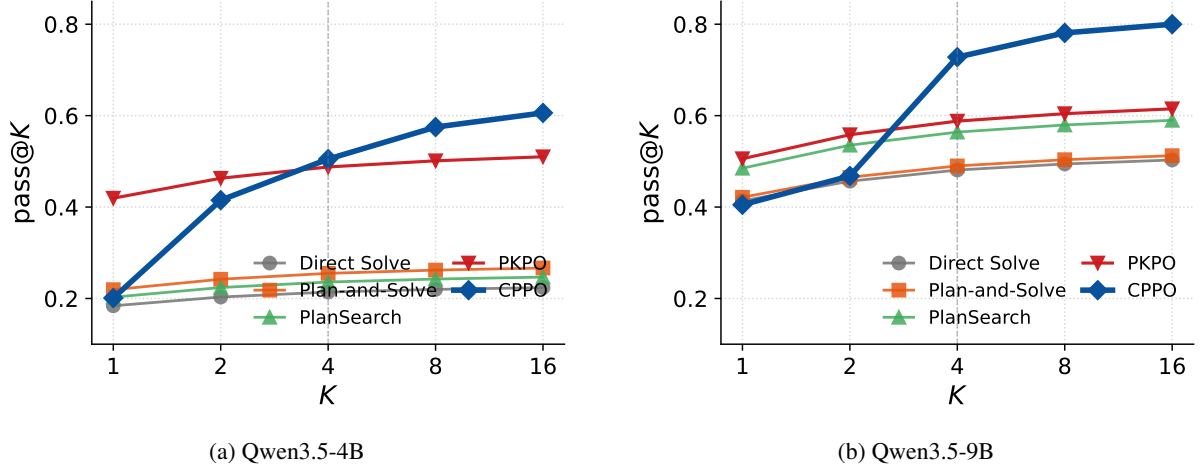


Figure 5: LiveCodeBench-v6 pass@ K across baselines for Qwen3.5-4B and Qwen3.5-9B.

of comparable size (Direct Solve 0.640/0.750 on E2B/E4B versus 0.420/0.515 on Qwen3.5-2B/4B), so the absolute CPPO margin over Direct Solve compresses (+0.108/ + 0.082 on E2B/E4B, versus +0.126/ + 0.279 on Qwen3.5-2B/4B) – the dominant headroom on APPS is in the strongest baseline rather than the base model. The LCBv6 pattern matches the Qwen3.5-2B sweep more directly.

P APPS Pass@ K Sweep

Figure 6 shows the full APPS pass@ K line sweep for all baselines across Qwen3.5-{2B, 4B, 9B}; Figure 7 gives the same data as a grouped-bar view at $K \in \{1, 2, 4, 8, 16\}$ for the 4B and 9B sizes.

Coordination unit and per- K_{solve} protocol. CPPO’s coordination unit is K_{tuple} : one planner rollout emits a single coordinated tuple of size K_{tuple} , and the solver attempts each strategy in the tuple. The Figure 2 sweep uses one CPPO planner trained at $K_{\text{tuple}}=4$ throughout (the same checkpoint as Table 1, with no K -specific retraining). For $K_{\text{solve}} \leq 4$ we take the first K_{solve} outputs of one 4-tuple rollout; for $K_{\text{solve}} > 4$ we pool $\lceil K_{\text{solve}}/4 \rceil$ independent 4-tuple rollouts. The solver outputs are pooled before computing pass@ K_{solve} . Coordination is preserved *within* each tuple, not across them: for $K_{\text{solve}} > 4$, pass@ K_{solve} measures pooled coordinated tuples versus iid samples at matched K_{solve} , not a single K_{solve} -way joint coordination. Table 22 gives the decomposition for each K_{solve} on the sweep of Figure 2.

Every method receives the same K_{solve} solver-attempt budget. Direct Solve draws K_{solve} in-

dependent base-model samples. Plan-and-Solve and PlanSearch draw K_{solve} samples through each method’s plan-selection procedure. PKPO and Pass@ K -only RL sample K_{solve} outputs from their trained policies. CPPO additionally consumes planner tokens at each tuple rollout; decoded-token accounting is reported in Appendix M.

Q Practical Rollout Recipe

A minimal CPPO run proceeds in a short pipeline, each stage building on the artifacts produced by the previous one.

We begin by assembling the planner-candidate pool: we sample from the current planner checkpoint under the same prompt that CPPO will later use. An offline LLM judge (Zheng et al., 2023) labels every candidate against a structured plan-validity rubric, and we balance the resulting train and validation splits by prompt and by binary label.

We then train a small planner reward model on this dataset with binary cross-entropy on the pass/fail target. Warm-up follows: holding the reward model frozen, we update only planner tokens under $R_{\text{warm}} = J_{\psi}(x, S)$ and stop when the reward-model pass rate plateaus. We integrate the frozen reward model into the planner-solver environment and run a brief smoke test under $R_{\text{plan}} = J_{\psi} \cdot R_{\text{out}}$, confirming that rewards are nonzero before we commit to a larger run.

The main run launches from this checkpoint and tracks the full diagnostic panel: outcome reward, plan-validity reward, the multiplicative planner reward, reward sparsity, and planner and solver quality, so that drift in any of them is caught early.

Because the planner’s distribution shifts as train-

Table 21: Pass@4 on APPS and LiveCodeBench-v6 for Gemma 4 E2B and E4B (Google DeepMind, 2026). LCBv6 Direct Solve cells are taken from the released Gemma 4 LCBv6 evaluations (Google DeepMind, 2026); all other cells are trained and evaluated under the protocol of §4.3 at the same $K=4$ solver-attempt budget.

Method	APPS		LiveCodeBench-v6	
	E2B	E4B	E2B	E4B
Direct Solve	0.640 ± 0.019	0.750 ± 0.018	$0.440^\dagger$	$0.520^\dagger$
Plan-and-Solve	0.660 ± 0.020	0.770 ± 0.019	0.452 ± 0.014	0.531 ± 0.017
PlanSearch	0.690 ± 0.022	0.790 ± 0.020	0.481 ± 0.014	0.560 ± 0.017
Pass@ K Training / RLVR	0.708 ± 0.018	0.798 ± 0.015	0.519 ± 0.019	0.598 ± 0.015
PKPO	0.725 ± 0.020	0.815 ± 0.017	0.521 ± 0.017	0.601 ± 0.015
CPPO	0.748 ± 0.021	0.832 ± 0.018	0.561 ± 0.019	0.629 ± 0.020

[†] Single-point estimate from (Google DeepMind, 2026); seed variance not reported.

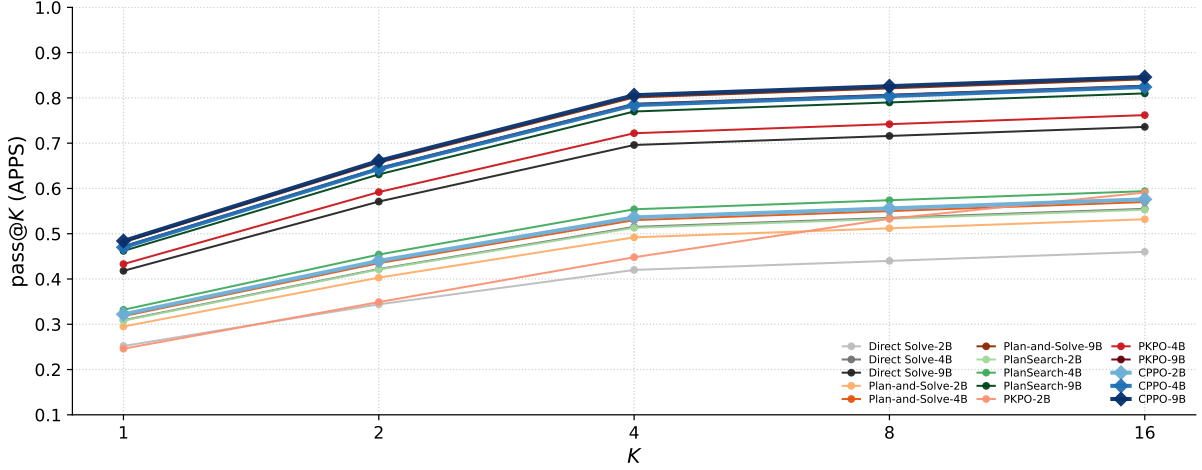


Figure 6: APPS pass@ K across Qwen3.5-2B, Qwen3.5-4B, and Qwen3.5-9B. Each method occupies a single color, with darker shades for larger models.

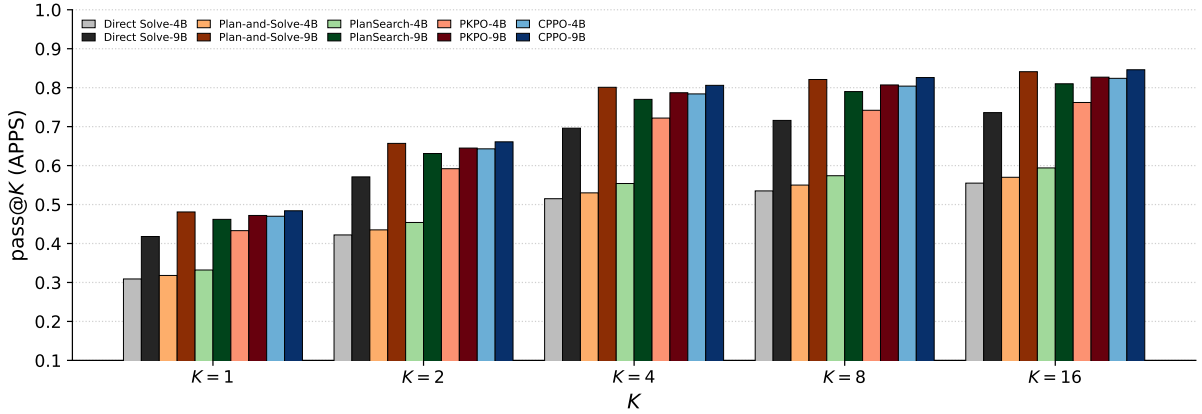


Figure 7: Grouped APPS pass@ K for Qwen3.5-4B and Qwen3.5-9B at $K \in \{1, 2, 4, 8, 16\}$. Each method occupies a single color across both sizes, with a lighter shade for 4B and a darker shade for 9B.

ing proceeds, we periodically refresh the reward-model dataset: we label outputs from the updated planner with the same judge, append them to the training pool, rebalance by prompt and pass/fail label, and resume finetuning the reward model from its latest checkpoint before CPPO continues.

R Baseline Implementation Details

This section details how the baselines reported in Table 1 are implemented. All methods use the same code-extraction, sandboxed execution, verifier, and pass@ K definition, and all are evaluated under the same K -attempt budget as CPPO. The inference-only baselines keep the base-model solver frozen; the trained baselines, PKPO and UpSkill, update

Table 22: CPPO rollout composition for the pass@ K sweep of Figure 2. The same $K_{\text{tuple}}=4$ planner is used at every K_{solve} : rows with $K_{\text{solve}} < 4$ truncate one 4-tuple to its first K_{solve} outputs; rows with $K_{\text{solve}} \geq 4$ pool $\lceil K_{\text{solve}}/4 \rceil$ independent 4-tuple rollouts. Coordination is preserved within each 4-tuple, not across tuples.

K_{solve}	CPPO rollout composition
1	first 1 of one 4-tuple
2	first 2 of one 4-tuple
4	one 4-tuple
8	two 4-tuples (4+4)
16	four 4-tuples (4+4+4+4)
20	five 4-tuples (4×5)
24	six 4-tuples (4×6)

the policy and then emit K answer samples at test time.

Training and selection protocol for trained baselines. PKPO and UpSkill use the same Qwen3.5 model sizes and the same CodeContests train pool as the CPPO rollout stage. Their optimization recipes are fixed on the CodeContests-train dev subset listed in Table 11, with no APPS, CodeContests-valid, or LCBv6 held-out examples used for hyperparameter choice, checkpoint selection, or early stopping. To keep the comparison controlled, we reuse the CPPO/PKPO optimizer and sampling settings where applicable: fp32 trainable weights, AdamW learning rate 5×10^{-7} , clip ratio $\epsilon=0.2$, KL coefficient 0.01, temperature 0.7, and maximum 2048 decoded tokens. Method-specific knobs are set to the paper defaults or to the pass@4 budget: PKPO uses $n=16$ and $k_{\text{opt}}=4$, while UpSkill uses $n_{\text{skills}}=4$ and $\alpha_{\text{mi}}=5.0$.

Plan-and-Solve (Wang et al., 2023a) as a prompting baseline.

1. The same base-model solver is given a prompt that asks it to first write a solution plan and then write the code.
2. For each problem we sample K full solution attempts under that prompt.
3. Each attempt passes through the same code extraction, sandboxed execution, and verifier pipeline.
4. Pass@ K is the fraction of problems on which at least one of the K attempts passes the tests.

This comparison isolates the effect of adding a planning prompt before code generation.

PlanSearch (Wang et al., 2024) as a planning baseline. We follow the released implementation, using eight candidate plans per problem and the same solver/verifier pipeline as the other baselines.

1. The base model first generates multiple candidate plans per problem; we use 8 candidates.
2. The candidates are selected and organized into plans usable for solving.
3. The same frozen base-model solver then generates code conditioned on each selected plan.
4. Each problem is evaluated under a budget of K solver attempts, with the same verifier and pass@ K definition as above.

This comparison isolates inference-time multi-plan search without planner training.

PKPO (Walder and Karkhanis, 2025) as a pass@ K RL baseline. PKPO transforms the per-sample reward vector $r \in \{0, 1\}^n$ over $n \geq k$ samples of the same problem into a low-variance unbiased estimator of the pass@ k gradient. We re-implement it on Qwen3.5-{2B, 4B, 9B} for code generation (the original paper uses Gemma-2 on math): each gradient step samples $n=16$ completions per problem, scores them with the same sandboxed verifier, and applies PKPO’s transform with $k_{\text{opt}}=4$ to obtain the REINFORCE-style update; no planner and no strategy tuple. At test time the trained policy emits $K=4$ independent answer samples and is evaluated under the same pass@ K verifier as CPPO. Training dynamics for the three PKPO runs are shown in Appendix S; we use them as a sanity check that the reimplementa-tion behaves as PKPO intends (variance-reduced, advantage-sparsifying updates) before treating it as a baseline.

UpSkill (Shah et al., 2026) as a diversity-oriented RL baseline. UpSkill is also a trained baseline rather than a frozen-solver prompt. It keeps the independent-answer policy class but conditions each sample on a latent skill prefix; at $K=4$, we draw one completion from each of four skill prefixes and score the resulting set with the same verifier. We adapt UpSkill to the same full-finetuning setting as PKPO and CPPO: the original rank-32 LoRA adapter is replaced by full-parameter training, while the MI weight $\alpha_{\text{mi}}=5.0$ and the remaining optimizer settings follow the protocol above. As with PKPO, the training dynamics in Appendix T confirm that the full-finetuning adap-

tation preserves UpSkill’s intended behavior—a mutual-information reward that rises alongside the correctness reward—before we use it as a baseline.

S PKPO Training Dynamics

Figure 8 tracks the PKPO update across 30 epochs for Qwen3.5-2B, 4B, and 9B, trained on the same CodeContests train pool as our other RL methods (see Table 3). Each step draws $n=16$ completions per problem at temperature 0.7, applies the SLOO-minus-one transform of Walder and Karkhanis (2025) with $k_{\text{opt}}=4$, and updates with AdamW at learning rate 5×10^{-7} , clip ratio $\epsilon=0.2$, and KL coefficient 0.01 against a frozen reference. The three sizes share these settings; only the base checkpoint differs.

The gradient and parameter-step norms decay as the policy converges, with 9B contracting fastest and 2B plateauing highest—larger models spread the same outcome signal over more parameters. The advantage nonzero rate is the fraction of problems still producing gradient: the SLOO-minus-one transform zeros it out whenever a problem is entirely solved or entirely failed, and larger models cross out of the all-fail regime sooner (≈ 0.55 at 2B versus ≈ 0.85 at 9B). The transformed-reward variance rises early as the policy first generates mixed pass/fail groups, plateaus through mid-training, and falls late as the policy converges on the solvable subset; the late drop is sharpest at 9B and most muted at 2B.

The 2B trace also matches the PKPO 2B rows of Table 8, where PKPO trails Tuple Planner SFT on both APPS and CodeContests: the elevated gradient plateau, the lowest advantage-nonzero rate of the three sizes, and the muted late variance drop together describe a sparse-signal regime in which most groups remain entirely failed and the SLOO-minus-one update has little to work with.

T UpSkill Training Dynamics

UpSkill shares PKPO’s policy class but adds a token-level mutual-information reward r_{mi} to the binary correctness reward r_{correct} . The MI term sets UpSkill apart: PKPO’s signal vanishes whenever a problem is entirely solved or entirely failed, while UpSkill’s MI term continues to produce a nonzero per-token gradient as long as the latent skill code distinguishes trajectories. The diagnostic panels therefore differ from Figure 8: we report the gradient L2 norm, the two reward components, and

the combined-reward variance.

Each step draws one completion per skill at temperature 0.7 under a fixed prompt prefix encoding the latent skill, with $n_{\text{skills}}=4$ (matched to the $K=4$ inference budget) and MI weight $\alpha_{\text{mi}}=5.0$. Optimizer settings mirror PKPO above; we replace the rank-32 LoRA adapter of the original UpSkill with full fine-tuning across all three sizes, leaving the rest of the hyperparameters unchanged. Parameter-step norm is omitted because the training logger did not record it.

The gradient L2 norm decays for every size, with the 9B trace smoother and the 2B trace visibly noisier. The correctness reward r_{correct} rises along a sigmoid whose asymptote scales with model size, consistent with stronger models reaching higher pass@4. The MI reward r_{mi} saturates early and stays high across all sizes: the latent-skill conditioning keeps trajectories distinguishable even after correctness plateaus, so the MI term can continue to supply token-level signal after correctness rewards become sparse. The combined-reward variance falls over training as the policy concentrates on a smaller set of higher-quality skills.

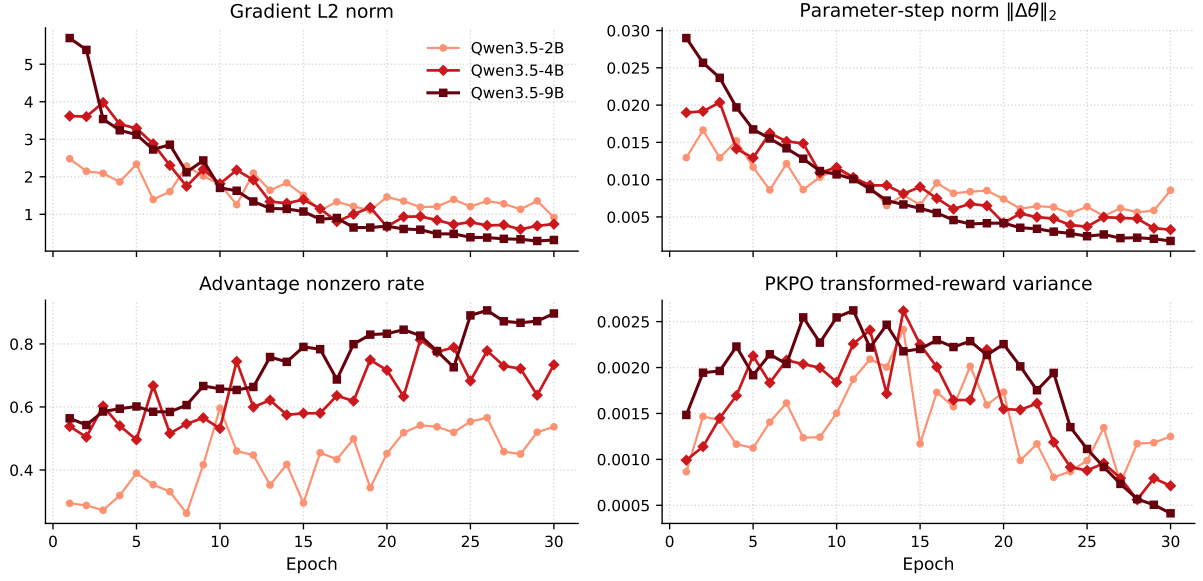


Figure 8: PKPO training dynamics for Qwen3.5-2B, 4B, and 9B over 30 epochs (light to dark red): gradient L2 norm, parameter-step norm $\|\Delta\theta\|_2$, advantage nonzero rate after group normalization, and variance of the PKPO-transformed reward.

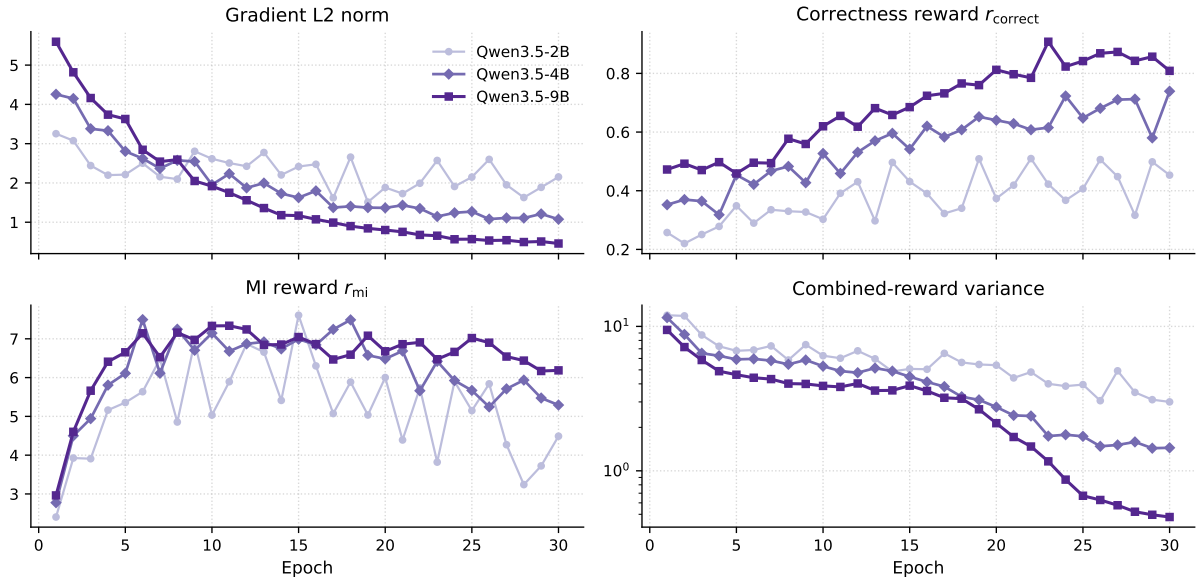


Figure 9: UpSkill training dynamics for Qwen3.5-2B, 4B, and 9B over 30 epochs (light to dark purple): gradient L2 norm, correctness reward r_{correct} , mutual-information reward r_{mi} , and combined-reward variance.