

TD-M(PC)²: Improving Temporal-Difference MPC Through Policy Constraint

Haotian Lin¹ · Pengcheng Wang² · Jeff Schneider¹ · Guanya Shi¹

¹ Carnegie Mellon University · ² University of California, Berkeley

Carnegie Mellon University UC Berkeley

MAIN CLAIM MPC exploration is useful; the bottleneck is bootstrapping the actor outside planner-data support. TD-M(PC)² constrains exploitation while leaving the planner free.

Abstract

Simplified from the paper abstract.

Plan-based MBRL combines sampling-based MPC with learned value and policy priors, but value learning is trained almost entirely on planner-collected data. We identify a structural mismatch between the planner policy used for exploration and the nominal policy evaluated by the value function, causing persistent value overestimation in high-dimensional tasks.

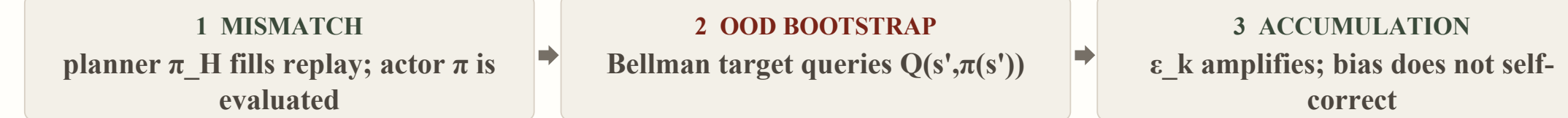
TD-M(PC)² adds a soft conservative policy update that keeps the learned actor close to replay-buffer behavior while preserving model-based exploration. Across HumanoidBench and DMControl, it reduces value bias and improves performance, with the largest gains on 61-DoF humanoid locomotion.

Problem π_H vs π mismatch	Fix actor-side $\beta \log \mu$	Evidence 2,159% bias + 37 tasks
---	---	---

Analysis

Why value bias persists in plan-based MBRL.

Mechanism: planner data are generated by π_H , but TD targets evaluate the nominal actor π . When π disagrees with π_H , $Q(s', \pi(s'))$ is queried outside replay support.



THEOREM 1: PLANNER PERFORMANCE BOUND

$$\limsup_{k \rightarrow \infty} |V^* - V^{\pi_{H,k}}| \leq \limsup_{k \rightarrow \infty} \frac{2}{1 - \gamma^H} \left[C(\epsilon_{m,k}, H, \gamma) + \frac{\epsilon_{p,k}}{2} + \frac{\gamma^H(1 + \gamma^2)}{(1 - \gamma)^2} \epsilon_k \right]$$

THEOREM 2: ERROR ACCUMULATION

$$\delta_k \leq \frac{1}{1 - \gamma^H} \left[2C(\epsilon_{m,k-1}, H, \gamma) + \epsilon_{p,k-1} + (1 + \gamma^H) \delta_{k-1} + \frac{2\gamma(1 + \gamma^{H-1})}{1 - \gamma} \epsilon_{k-1} \right]$$

DISTRIBUTION SHIFT LOWER BOUND

$$D_{TV}(p^\pi(s, a) \| p^{\pi'}(s, a)) \geq \frac{1 - \gamma}{2R_{\max}} |J^\pi - J^{\pi'}|$$

Interpretation. H-step planning reduces dependence on terminal value error, but the error itself grows when actor queries leave planner-data support.

Empirical signature. Value bias scales with action dimension: Dog-trot reaches 231%, h1hand-run reaches 2,159%, and h1hand-slide reaches 746% under vanilla TD-MPC2.

Method

Conservative exploitation; unchanged model-based exploration.

TD-M(PC)² adopts a soft state-conditional divergence constraint to the replay behavior policy μ , implemented using the stored planner Gaussian. The method changes the actor update, not the planner.

CONSTRAINED POLICY IMPROVEMENT

$$\pi_{k+1} \in \arg \max_{\pi} \mathbb{E}_{a \sim d_{\pi}} [\mathbb{E}_{a \sim \pi(\cdot|s)} Q_k(s, a) - \beta D(\pi(\cdot|s) \| \mu(\cdot|s))]$$

IMPLEMENTED ACTOR LOSS

$$\mathcal{L}_{\pi} = -\mathbb{E}_{s, \mu \sim B, a \sim \pi} \left[\frac{Q(s, a)}{S_q} + \frac{\beta \log \mu(a|s)}{S_q} - \alpha \log \pi(a|s) \right]$$

Store in replay planner mean/std with each transition	No density model reuse $\mu \approx \pi_H$; <10 LoC
--	---

Actor-side only MPC search, model, critic target unchanged	Optimistic conservatism entropy term + β threshold via S_q
---	---

LOSS SCALE

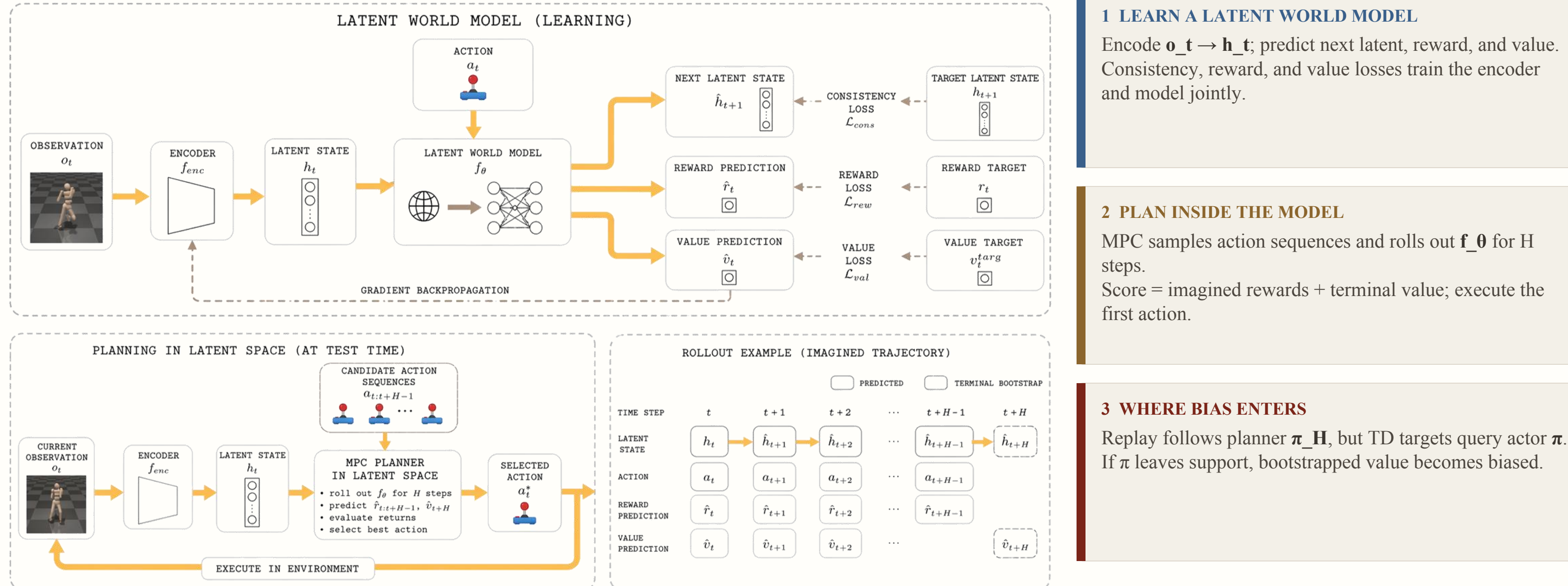
$$S_q = \text{EMA}(\max\{\text{Per}(Q, 95) - \text{Per}(Q, 5), 1\}, \xi)$$

Why it works. The constraint is applied only during actor improvement. The replay planner still explores with MPC, while the actor used in TD targets is discouraged from querying high-value unsupported actions. This attacks OOD bootstrapping without adding a density model or changing the world model, critic target, or planner objective.

Constrain exploitation; preserve exploration.

Architecture

TD-MPC latent world model learning, test-time MPC planning, and imagined rollouts.

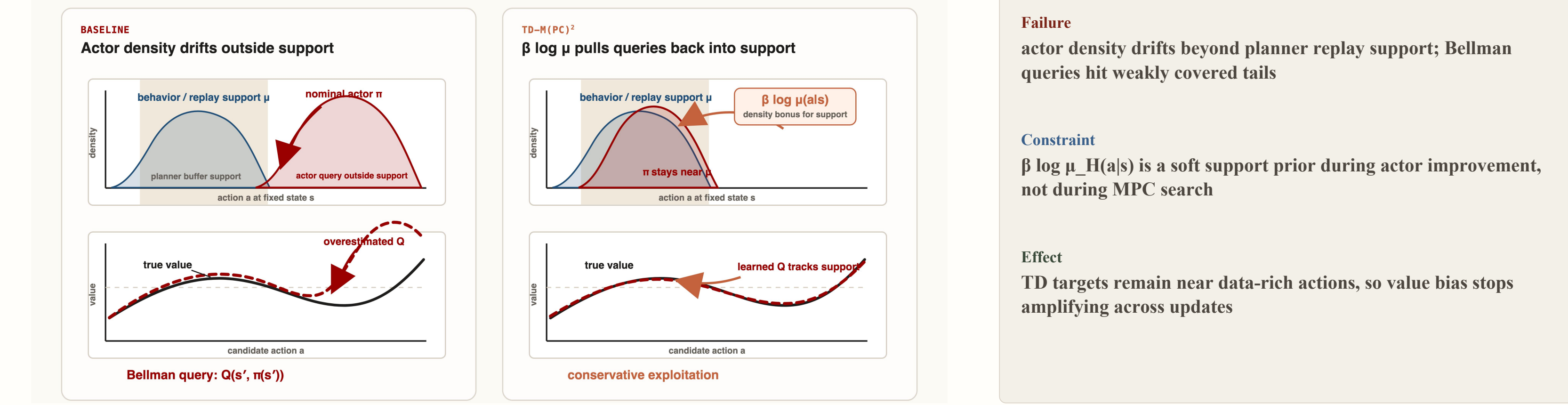


TD-M(PC)² intervention: keep the MBRL stack unchanged; constrain actor-side exploitation.

Stored support $\mu_H(a s) = \mathcal{N}(a; \hat{a}_H(s), \Sigma_H(s))$ store planner Gaussian parameters with replay	Actor update $Q_{\theta}(s, a) + \beta \log \mu_H(a s)$ penalize unsupported high-Q actions	Unchanged planner $f_{\theta}, r_{\theta}, Q_{\theta}, \pi_H$ unchanged MPC search and model learning stay intact
--	--	--

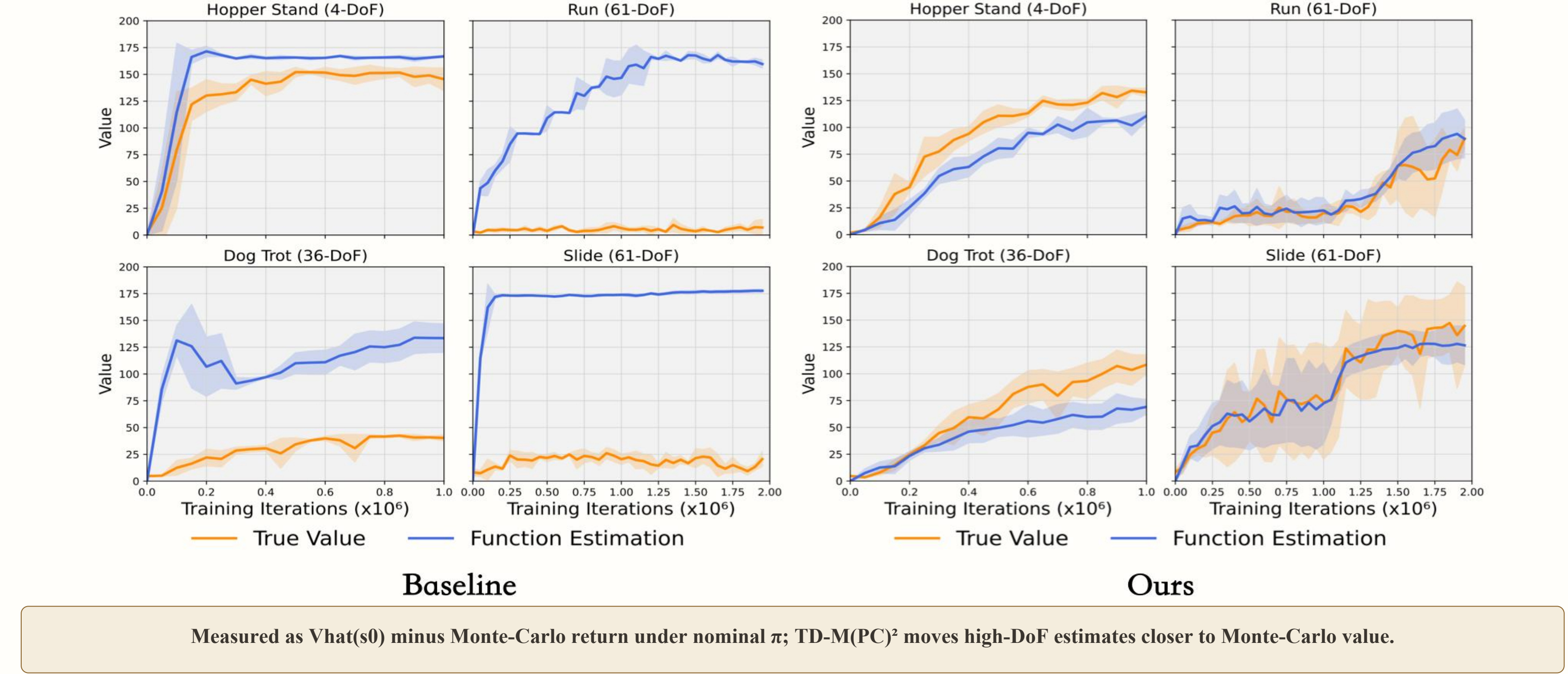
Mitigate Value Overestimation

$\beta \log \mu$ pulls actor queries back into replay support.



Value-bias diagnostic

Initial-state Monte-Carlo return under nominal actor versus critic estimate.



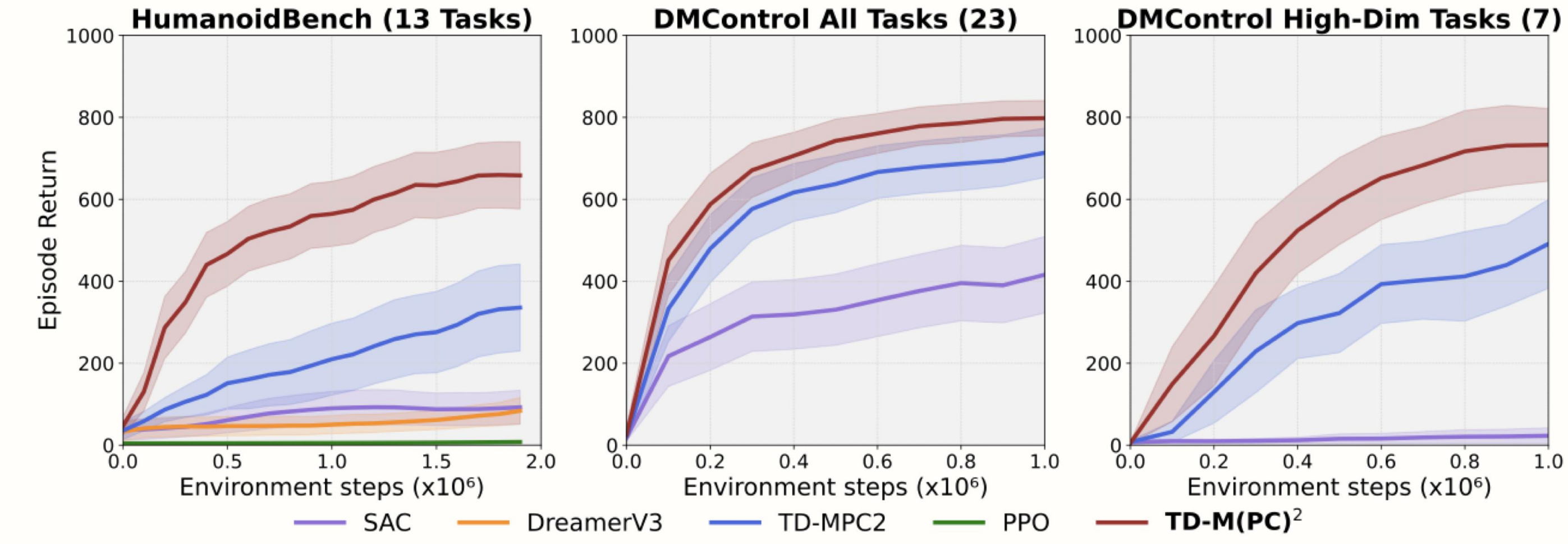
37
evaluation tasks
14 HB + 23 DMC

>100%
HB aggregate gain
vs TD-MPC2; Reach excluded

<10 LoC
policy-side change
same actor architecture

Average performance

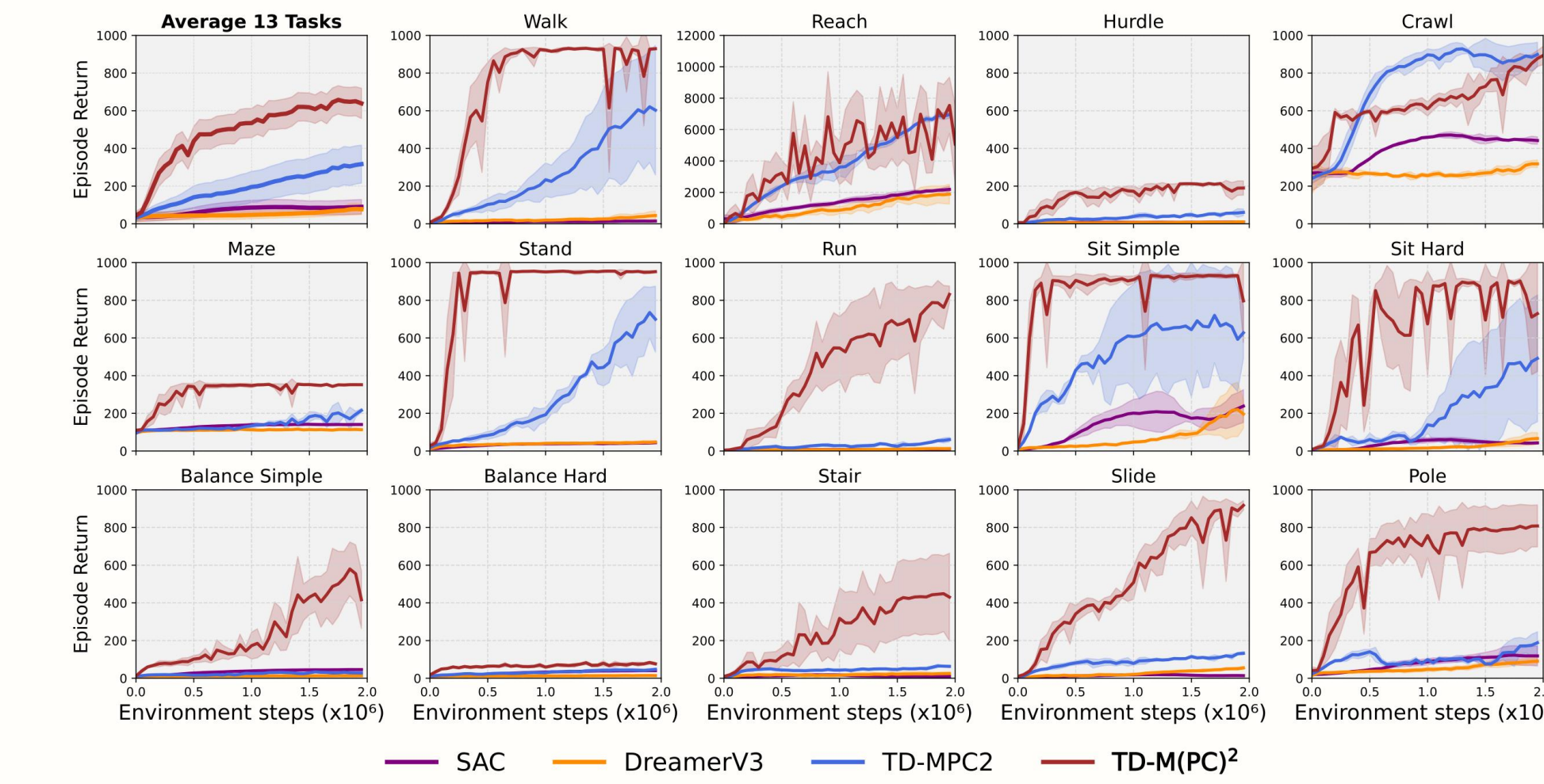
Main-paper aggregate across HumanoidBench and DMControl.



13 HumanoidBench avg	23 DMControl tasks	7 high-dim DMC	Takeaway. All aggregate suites improve; largest gains are on high-dimensional control.
--------------------------------	------------------------------	--------------------------	---

HumanoidBench per-task curves

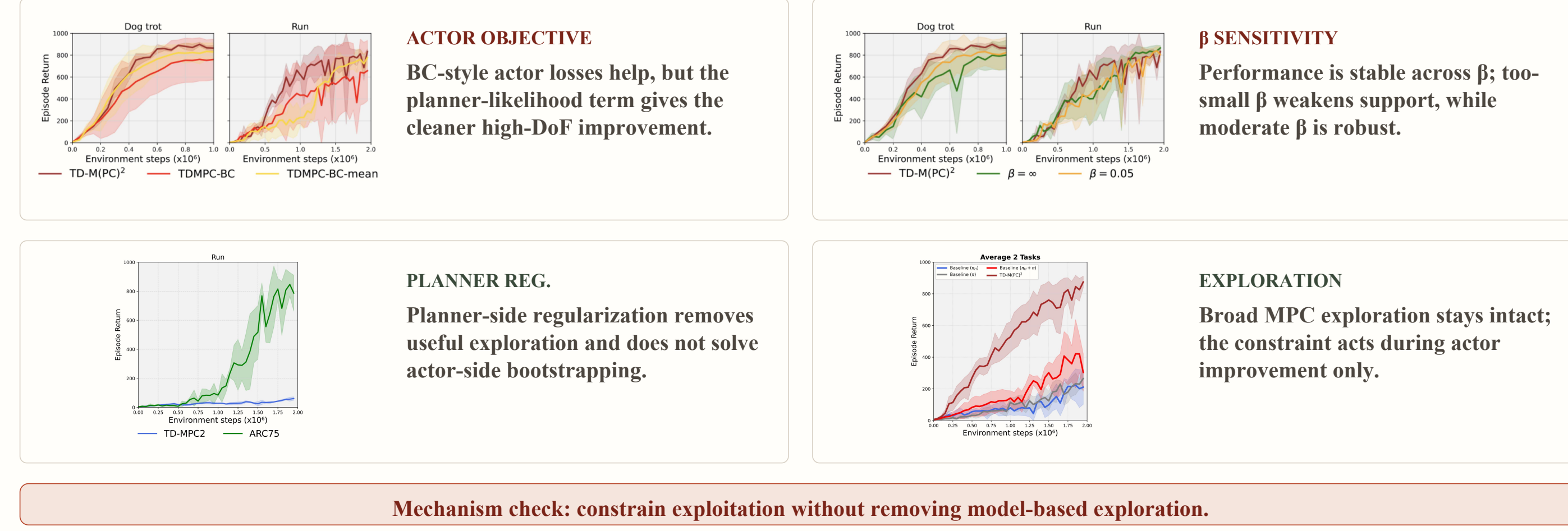
Complete h1hand-v0 locomotion results.



Full HB evidence. The gain is not a single-task artifact.
Hard 61-DoF tasks: run, slide, stair, pole.
Reach excluded due to reward scale. Pattern: support matters most when value errors drive large whole-body motions.

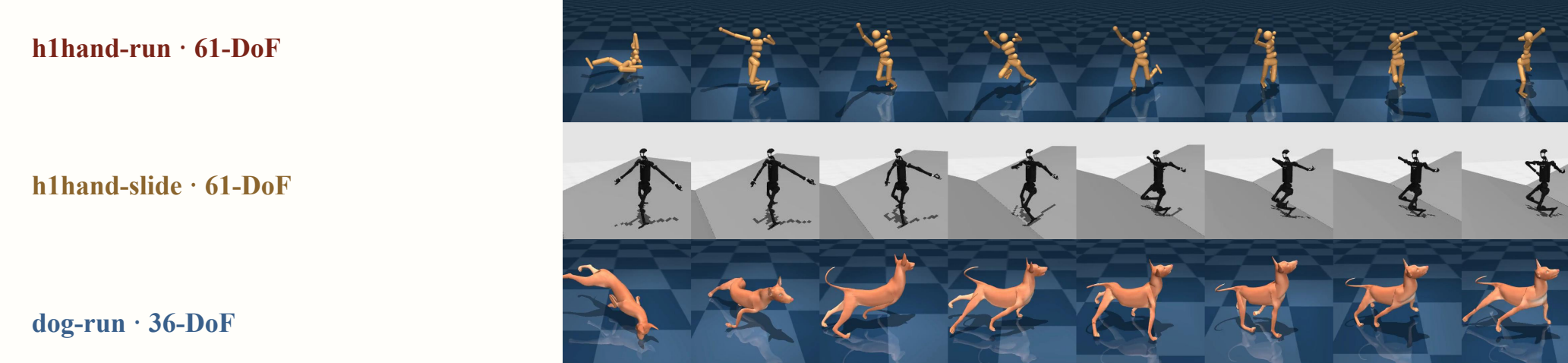
Ablations

Constraint mechanism, sensitivity, and planner-side checks.



Qualitative rollouts

Stable whole-body behavior on high-dimensional tasks.



Paper, code, videos, and reproducibility material
https://darthutopian.github.io/tdmpe_square/