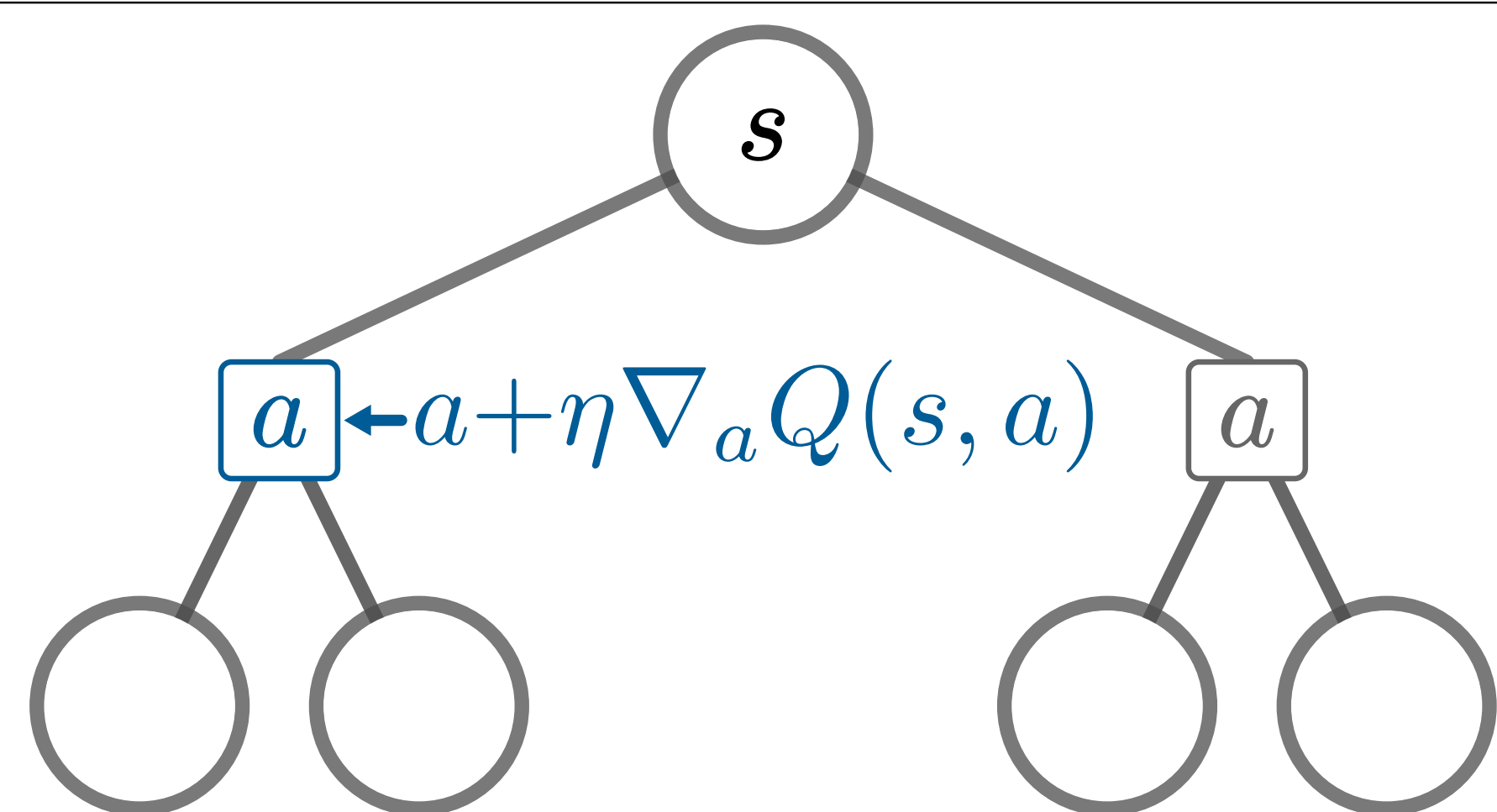


AGMCTS (Action-Gradient MCTS) combines *global action sampling*, *local refinement* and *consistent value estimation*

1. Refine Actions Inside Tree Search



MCTS samples distinct action branches, each carrying a $Q(s, a)$ estimate. AGMCTS updates selected actions with gradient information $\nabla_a Q$.

Action Score Gradient: MDP Case

Samples may come from an earlier action through a proposal q :

$$\nabla_a Q_t^\pi(s, a) = \mathbb{E}_{s' \sim q} \left[\rho_a(s') \left(\nabla_a \log p_T(s' | s, a) \cdot (r(s, a, s') + \gamma V_{t+1}^\pi(s') - V(s)) + \nabla_a r(s, a, s') \right) \right],$$

$$\rho_a(s') = p_T(s' | s, a) / q(s').$$

Reuse correction (Reweight earlier successors) **Transition score** (Favor high-return successors) **Direct reward** (Immediate reward gradient)

Extending to POMDPs via Particle-Belief MDPs

The exact density $p(b' | b, a)$ is intractable; we expand through the propagated belief b'^- and observation o' :

$$p(b' | b, a) = \int p(b' | b'^-, o') p(o' | b'^-) p(b'^- | b, a) db'^- do'.$$

For an ordered particle belief $\bar{b} = ((s^j, \lambda^j))_{j=1}^J$, it is tractable:

$$p(\bar{b}'^- | \bar{b}, a) = \prod_{j=1}^J p_T(s'^{-j} | s^j, a).$$

MDP: $p_T(s' | s, a) \implies$ POMDP: $\prod_{j=1}^J p_T(s'^{-j} | s^j, a)$
 The same action-score formula applies with adjusted likelihood

2. Reuse the Tree After Action Optimization

Each successor s'^i remembers the proposal action a_{prop}^i under which it was sampled. Let $N_i := n(s'^i) + 1$. The MIS Tree estimates future values by

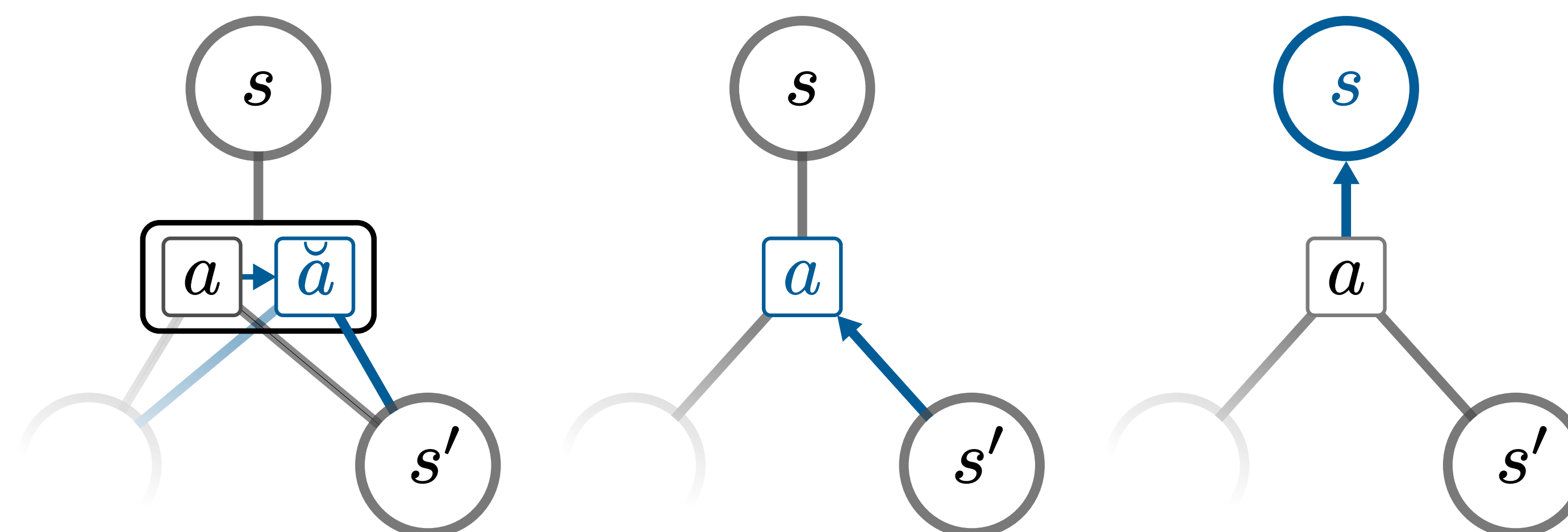
$$V_f(s, a) = \sum_i N_i \rho_a^i(s'^i) V(s'^i) / \eta(s, a), \quad \eta(s, a) = \sum_i N_i \rho_a^i(s'^i).$$

MDP ratio

$$\rho_a^i(s') = p_T(s' | s, a) / p_T(s' | s, a_{\text{prop}}^i).$$

POMDP ratio

$$\rho_a^i(\bar{b}'^-) = p(\bar{b}'^- | \bar{b}, a) / p(\bar{b}'^- | \bar{b}, a_{\text{prop}}^i).$$



1. Action Update

$a \rightarrow \tilde{a}$:

Update ratios in $O(|\mathcal{S}_a|)$

2. Action Backprop

$V(s') \rightarrow \tilde{V}(s')$:

Update $Q(s, a)$ in $O(1)$

3. State Backprop

$Q(s, a) \rightarrow \tilde{Q}(s, a)$:

Update $V(s)$ in $O(1)$

Value estimates are defined through $V(s) = \sum_a n(s, a) Q(s, a) / \sum_a n(s, a)$.

Transition Densities from Smooth Generative Models

Suppose $s' = F(s, a, \xi)$ is locally Lipschitz, with $\text{rank}(D_\xi F) = n_\xi$ almost everywhere. For noise $\xi \sim p_\xi(\cdot | s, a)$, the Area Formula gives the density induced on the successor manifold:

$$p_T(s' | s, a) = \sum_{\xi^*: F(s, a, \xi^*) = s'} p_\xi(\xi^* | s, a) / J_{n_\xi} F(s, a, \xi^*),$$

$$J_{n_\xi} F = \sqrt{\det((D_\xi F)^\top D_\xi F)}.$$

Noise before the simulator

$$h(s, a, \xi) = (\tilde{s}, \tilde{a}), \quad s' = f(\tilde{s}, \tilde{a})$$

$$D_\xi F = D_{(\tilde{s}, \tilde{a})} f \cdot \partial h / \partial \xi.$$

Unique noise inverse: reuse the cached simulator Jacobian.

Noise after the simulator

$$s' = g(h(s, a), \xi)$$

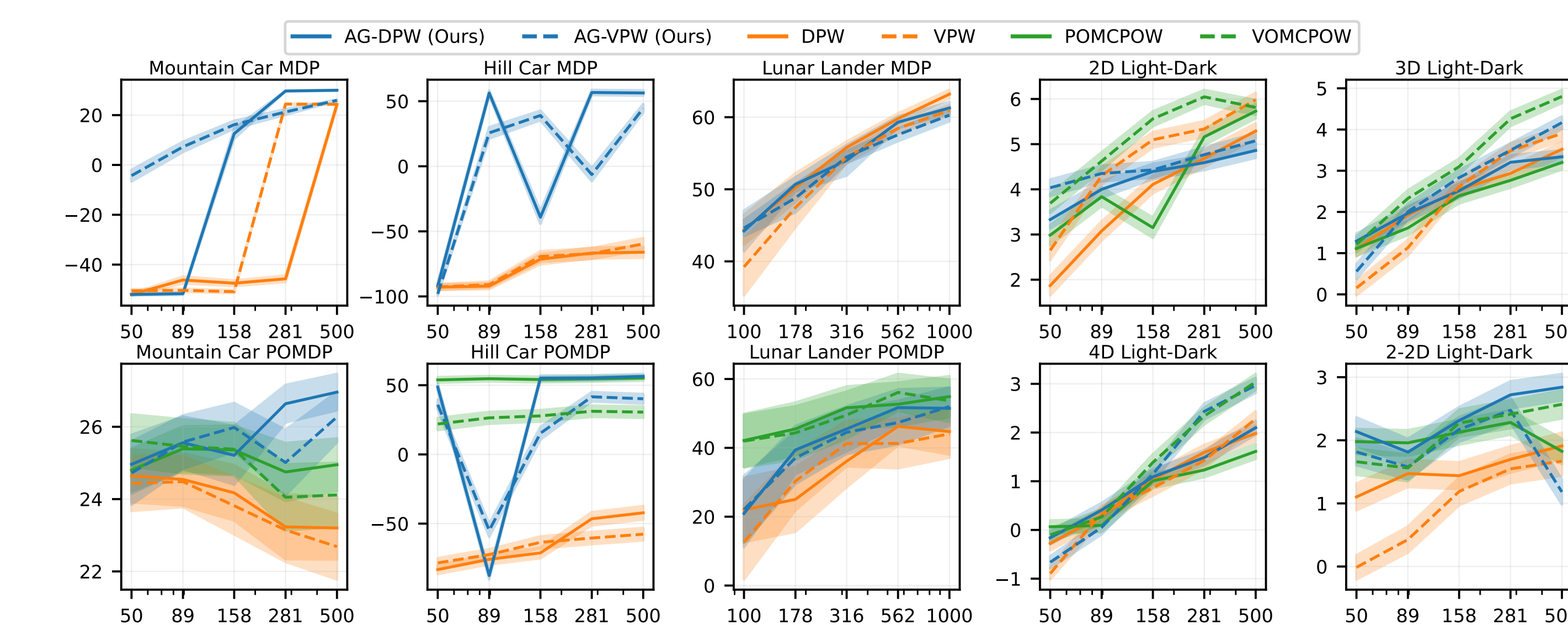
$$D_\xi F = \partial g / \partial \xi|_{h(s, a), \xi^*}.$$

Solvable noise preimages: one forward simulator evaluation.

3. Does Refinement Improve Planning?

We compare Double Progressive Widening (DPW) and Voronoi Progressive Widening (VPW) to their Action-Gradient (AG-DPW/AG-VPW) variants. In POMDPs, we compare to POMCPOW and VOMCPOW as well.

We evaluate in 3 MDPs and 7 POMDPs with continuous states, actions and observations, exhibiting stochastic transition and observation models.



Average return over 1000 shared seeds; shading shows ± 2 standard errors. The x-axis is the simulation budget per planning session. AG-DPW and AG-VPW are AGMCTS variants; PFT denotes the particle-filter-tree POMDP adaptation.

Empirical Bottom Lines

6/10

scenarios: at least one AG variant preferable to its matched baseline

10 / 7 / 3

improve / neither / underperform in matched AG comparisons

5/7

POMDP scenarios: AGMCTS best or statistically tied

- AG-preferred scenarios: all Mountain/Hill Car settings, 4D and two-agent Light-Dark
- Largest gains: Mountain/Hill Car; Light-Dark and Lunar Lander show more mixed outcomes

Runtime Trade-Off

10–35×

single-threaded per-plan overhead in MDPs

2–20×

single-threaded per-plan overhead in POMDPs

Practical trade-off: whether improved action quality per planner query is worth the additional density/gradient overhead within latency budget

Takeaway

Best fit: small action changes strongly affect states or rewards, like in long-horizon Mountain Car or continuous-reward Light-Dark

Requires: transition-density access and more hyperparameter tuning



Scan for paper