

Actor-Critic with a Baseline

The Advantage

Reinforcement Learning

Sai Sampath Kedari

sampath@umich.edu

Contents

1	Actor-Critic with a Baseline: The Advantage	2
1.1	Starting Point from QAC	2
1.2	Baseline Policy-Gradient Identity	2
1.3	Value Baseline and Advantage	3
1.4	From Two Critics to One	4
1.5	Expected SARSA Bellman Form	4
1.6	TD Error as One-Sample Advantage	5
1.7	Exact Identity vs Approximation	5
1.8	One-Critic Actor-Critic	6
1.9	Bias and Variance, Layer by Layer	6
1.10	Algorithm: One-Step Actor-Critic	6
1.11	The Ladder	7
1.12	What We Bought	7
1.13	Scope	7
A	Time-Marginal Form	7
B	Estimator Naming Ladder	8

1 Actor-Critic with a Baseline: The Advantage

1.1 Starting Point from QAC

Report 15 carried the exact policy gradient down to a per-step oracle estimator [report 15]:

Key Result

$$\hat{g}_t^{Q,\text{oracle}}(\theta) = \gamma^t Q^{\pi_\theta}(S_t, A_t) \nabla_\theta \log \pi_\theta(A_t | S_t).$$

This report changes only the scalar multiplier of the score:

$$Q^{\pi_\theta}(S_t, A_t) \rightsquigarrow Q^{\pi_\theta}(S_t, A_t) - b(S_t).$$

A state-only baseline $b = b_\theta(S)$ is allowed; $b(S, A)$ is not generally allowed. A baseline never changes the mean; a good baseline reduces variance. The whole report is the search for the multiplier that keeps the gradient exact while shrinking its variance.

The γ^t -weighted time-marginal machinery that produced this estimator is not repeated; it is identical to Report 15 with the multiplier carried through, and is recorded in Appendix A.

Everything downstream of the policy gradient theorem is a contest over one scalar: the multiplier on $\nabla_\theta \log \pi_\theta(A_t | S_t)$. Report 15 used Q^{π_θ} . Its flaw is not bias but variance. $Q^{\pi_\theta}(S_t, A_t)$ carries the full return magnitude, most of which is shared across all actions at S_t and says nothing about which action to prefer. This report subtracts that shared, action-independent part.

1.2 Baseline Policy-Gradient Identity

Theorem 1.1 (Baseline Variance Reduction). *For any state-only baseline $b(S)$, i.e. a function of the state but not the action,*

$$\begin{aligned} \nabla_\theta J(\theta) &= \frac{1}{1-\gamma} \mathbb{E}_{S \sim d^{\pi_\theta, \rho_0}} \left[\mathbb{E}_{A \sim \pi_\theta(\cdot | S)} \left[Q^{\pi_\theta}(S, A) \nabla_\theta \log \pi_\theta(A | S) \right] \right] \\ &= \frac{1}{1-\gamma} \mathbb{E}_{S \sim d^{\pi_\theta, \rho_0}} \left[\mathbb{E}_{A \sim \pi_\theta(\cdot | S)} \left[(Q^{\pi_\theta}(S, A) - b(S)) \nabla_\theta \log \pi_\theta(A | S) \right] \right]. \end{aligned}$$

Proof. It suffices to show the baseline term contributes nothing:

$$\mathbb{E}_{A \sim \pi_\theta(\cdot | s)} \left[b(s) \nabla_\theta \log \pi_\theta(A | s) \right] = 0 \quad \forall s.$$

Expanding the inner action expectation at a fixed state s ,

$$\begin{aligned}
\mathbb{E}_{A \sim \pi_\theta(\cdot | s)} [b(s) \nabla_\theta \log \pi_\theta(A | s)] &= \sum_a \pi_\theta(a | s) b(s) \nabla_\theta \log \pi_\theta(a | s) && \text{; definition of } \mathbb{E}_A \\
&= b(s) \sum_a \pi_\theta(a | s) \frac{\nabla_\theta \pi_\theta(a | s)}{\pi_\theta(a | s)} && \text{; } \nabla_\theta \log \pi = \nabla_\theta \pi / \pi \\
&= b(s) \sum_a \nabla_\theta \pi_\theta(a | s) && \text{; cancel } \pi_\theta(a | s) \\
&= b(s) \nabla_\theta \left(\sum_a \pi_\theta(a | s) \right) && \text{; gradient out of finite sum} \\
&= b(s) \nabla_\theta(1) = 0. && \text{; } \sum_a \pi_\theta(a | s) = 1
\end{aligned}$$

Subtracting a term of zero mean from the multiplier leaves the outer expectation unchanged, which is the second form. $\square$

Key Result

$$\nabla_\theta J(\theta) = \frac{1}{1-\gamma} \mathbb{E}_{S \sim d^{\pi_\theta, \rho_0}} \left[\mathbb{E}_{A \sim \pi_\theta(\cdot | S)} \left[(Q^{\pi_\theta}(S, A) - b(S)) \nabla_\theta \log \pi_\theta(A | S) \right] \right].$$

The step $\sum_a \nabla_\theta \pi_\theta(a | s) = \nabla_\theta(1) = 0$ is the same expected-score-zero identity used in the policy-gradient preliminaries: the score has mean zero under π_θ , so any state-only weight $b(S)$ washes out.

1.3 Value Baseline and Advantage

Choose

$$b(S) = V^{\pi_\theta}(S).$$

Then

$$\begin{aligned}
Q^{\pi_\theta}(S, A) - b(S) &= Q^{\pi_\theta}(S, A) - V^{\pi_\theta}(S) && \text{; value baseline} \\
&= A^{\pi_\theta}(S, A). && \text{; definition of advantage}
\end{aligned}$$

Therefore,

Key Result

$$\nabla_\theta J(\theta) = \frac{1}{1-\gamma} \mathbb{E}_{S \sim d^{\pi_\theta, \rho_0}} \left[\mathbb{E}_{A \sim \pi_\theta(\cdot | S)} \left[A^{\pi_\theta}(S, A) \nabla_\theta \log \pi_\theta(A | S) \right] \right].$$

Two expectations of the same return, one conditioned on the action, one averaging it out:

$$A^{\pi_\theta}(S_t, A_t) = \underbrace{\mathbb{E}^{\pi_\theta}[G_t | S_t, A_t]}_{Q^{\pi_\theta}(S_t, A_t)} - \underbrace{\mathbb{E}^{\pi_\theta}[G_t | S_t]}_{V^{\pi_\theta}(S_t)}.$$

The advantage measures how much taking A_t shifts the expected return relative to the policy's average at S_t . The score direction is now signed by this comparison, not by the raw magnitude of Q^{π_θ} .

Equivalently, the advantage is Q^{π_θ} centered to zero mean under π_θ :

$$\mathbb{E}_{A \sim \pi_\theta(\cdot | S)}[A^{\pi_\theta}(S, A)] = \mathbb{E}_{A \sim \pi_\theta(\cdot | S)}[Q^{\pi_\theta}(S, A)] - V^{\pi_\theta}(S) = 0.$$

This is why V^{π_θ} is the *natural* baseline, not merely *a* baseline: it is the state-only offset that removes exactly the action-independent part, leaving a mean-zero multiplier.

Carrying the new multiplier through the same per-step form as Report 15 gives the advantage oracle estimator:

Key Result

$$\begin{aligned} \hat{g}_t^{A, \text{oracle}}(\theta) &= \gamma^t A^{\pi_\theta}(S_t, A_t) \nabla_\theta \log \pi_\theta(A_t | S_t) \\ &= \gamma^t [Q^{\pi_\theta}(S_t, A_t) - V^{\pi_\theta}(S_t)] \nabla_\theta \log \pi_\theta(A_t | S_t). \end{aligned}$$

1.4 From Two Critics to One

The advantage needs two learned functions,

$$A^{\pi_\theta}(S_t, A_t) = Q^{\pi_\theta}(S_t, A_t) - V^{\pi_\theta}(S_t) \rightsquigarrow \hat{q}(S_t, A_t, w^q) - \hat{v}(S_t, w^v),$$

two critics and two parameter vectors. The next two subsections collapse this to a single V -critic: Q^{π_θ} is rewritten through its Bellman equation so that the whole advantage is estimated by one TD error.

1.5 Expected SARSA Bellman Form

Key Result

$$Q^{\pi_\theta}(s, a) = \mathbb{E}^{\pi_\theta}[R_{t+1} + \gamma V^{\pi_\theta}(S_{t+1}) | S_t = s, A_t = a].$$

This is the Expected SARSA Bellman form for Q^{π_θ} .

The next action is already averaged inside V^{π_θ} :

$$V^{\pi_\theta}(s') = \mathbb{E}_{A' \sim \pi_\theta(\cdot | s')} [Q^{\pi_\theta}(s', A')].$$

Two targets for the same Q^{π_θ} :

$$\begin{aligned} R_{t+1} + \gamma Q^{\pi_\theta}(S_{t+1}, A_{t+1}) & \quad \text{SARSA target: samples } A_{t+1}, \\ R_{t+1} + \gamma V^{\pi_\theta}(S_{t+1}) & \quad \text{Expected SARSA target: no } A_{t+1} \text{ sample.} \end{aligned}$$

In QAC the SARSA bootstrap samples the next action A_{t+1} to form the target. The Expected SARSA target bootstraps with $V^{\pi_\theta}(S_{t+1})$, the exact policy average over A_{t+1} : the next action is integrated out analytically rather than sampled. This is Rao-Blackwellization, replacing a sample by its conditional expectation can only lower variance, never raise it. One fewer sampled variable in the target, and since the advantage never needs Q^{π_θ} alone but only its deviation from V^{π_θ} , a single V -critic suffices. That is the estimation advantage of this report over QAC.

1.6 TD Error as One-Sample Advantage

Define the true TD error:

$$\delta_t^\pi := R_{t+1} + \gamma V^{\pi_\theta}(S_{t+1}) - V^{\pi_\theta}(S_t).$$

Theorem 1.2 (TD Error is a One-Sample Advantage). *Under π_θ ,*

$$\mathbb{E}^{\pi_\theta} [\delta_t^\pi \mid S_t = s, A_t = a] = A^{\pi_\theta}(s, a).$$

Proof.

$$\begin{aligned} \mathbb{E}^{\pi_\theta} [\delta_t^\pi \mid s, a] &= \mathbb{E}^{\pi_\theta} [R_{t+1} + \gamma V^{\pi_\theta}(S_{t+1}) - V^{\pi_\theta}(S_t) \mid S_t = s, A_t = a] && \text{; definition} \\ &= \mathbb{E}^{\pi_\theta} [R_{t+1} + \gamma V^{\pi_\theta}(S_{t+1}) \mid S_t = s, A_t = a] - V^{\pi_\theta}(s) && \text{; } S_t = s \\ &= Q^{\pi_\theta}(s, a) - V^{\pi_\theta}(s) && \text{; Expected SARSA Bellman form} \\ &= A^{\pi_\theta}(s, a). && \text{; definition of advantage} \end{aligned}$$

□

The advantage never had to be built from two critics: one sample of the TD error is already an unbiased sample of it. This gives the oracle estimator

$$\hat{g}_t^{\delta, \text{oracle}}(\theta) = \gamma^t \delta_t^\pi \nabla_\theta \log \pi_\theta(A_t \mid S_t).$$

1.7 Exact Identity vs Approximation

The exact identity

$$\mathbb{E}^{\pi_\theta} [\delta_t^\pi \mid S_t, A_t] = A^{\pi_\theta}(S_t, A_t)$$

holds for

$$\delta_t^\pi = R_{t+1} + \gamma V^{\pi_\theta}(S_{t+1}) - V^{\pi_\theta}(S_t).$$

In practice,

$$\delta_t = R_{t+1} + \gamma \hat{v}(S_{t+1}, w_t) - \hat{v}(S_t, w_t).$$

Then

$$\mathbb{E}^{\pi_\theta} [\delta_t \mid S_t, A_t] \approx A^{\pi_\theta}(S_t, A_t), \quad \text{with equality only if } \hat{v}(\cdot, w_t) = V^{\pi_\theta}(\cdot).$$

1.8 One-Critic Actor-Critic

$$V^{\pi_\theta}(s) \rightsquigarrow \hat{v}(s, w).$$

Key Result

$$\begin{aligned}\delta_t &= R_{t+1} + \gamma \hat{v}(S_{t+1}, w_t) - \hat{v}(S_t, w_t), \\ w_{t+1} &= w_t + \alpha_t^w \delta_t \nabla_w \hat{v}(S_t, w_t), \\ \theta_{t+1} &= \theta_t + \alpha_t^\theta \gamma^t \delta_t \nabla_\theta \log \pi_\theta(A_t | S_t)|_{\theta=\theta_t}.\end{aligned}$$

1.9 Bias and Variance, Layer by Layer

	Object	Bias	Variance
A	$\hat{g}_t^{Q, \text{oracle}}$	none, oracle Q^π	high
B	$Q^\pi(S_t, A_t) - b(S_t)$	none	changed
C	$Q^\pi(S_t, A_t) - V^\pi(S_t) = A^\pi(S_t, A_t)$	none	usually lower
D	δ_t^π	none as an advantage sample	TD target sampling
E	$\delta_t = R_{t+1} + \gamma \hat{v}(S_{t+1}, w_t) - \hat{v}(S_t, w_t)$	critic approximation/tracking	TD target sampling

1.10 Algorithm: One-Step Actor-Critic

Algorithm 1: One-Step Actor-Critic, Rollout Form

Input: a differentiable policy parameterization $\pi_\theta(a | s)$

Input: a differentiable state-value parameterization $\hat{v}(s, w)$

Input: step sizes $\alpha_t^\theta, \alpha_t^w > 0$; discount γ

initialize policy parameter $\theta_0 \in \mathbb{R}^{d'}$ and state-value weights $w_0 \in \mathbb{R}^d$

repeat

 sample/reset $S_0 \sim \rho_0$

$S_t \leftarrow S_0$

for $t = 0, 1, \dots, T - 1$ *or until* S_t *terminal* **do**

 Sample $A_t \sim \pi_{\theta_t}(\cdot | S_t)$

 Observe $S_{t+1}, R_{t+1} \sim p(\cdot, \cdot | S_t, A_t)$

// environment dynamics

$\delta_t \leftarrow R_{t+1} + \gamma \hat{v}(S_{t+1}, w_t) - \hat{v}(S_t, w_t)$

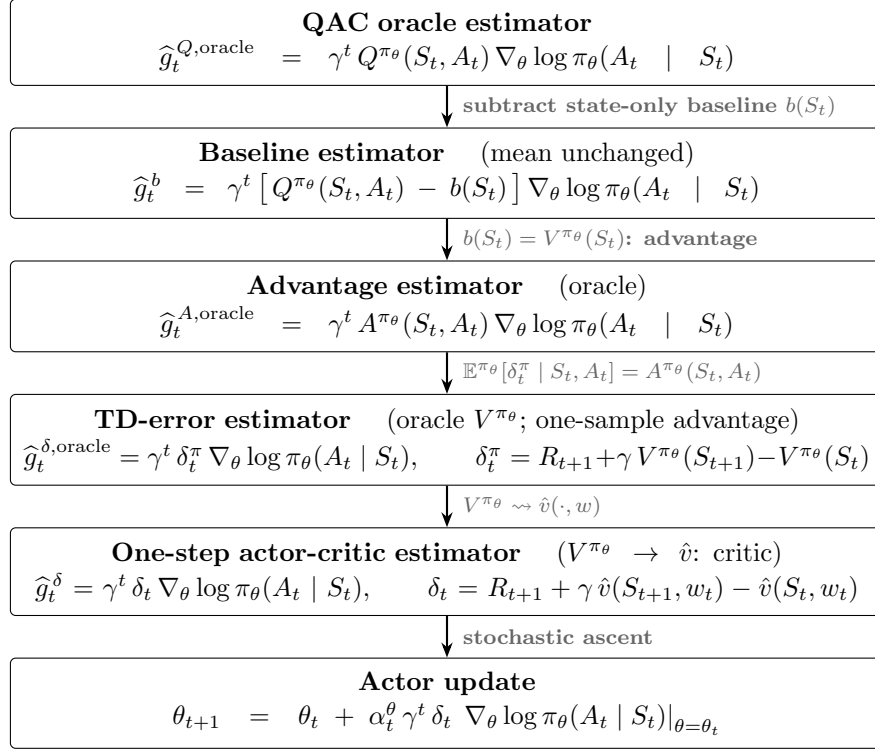
$\theta_{t+1} \leftarrow \theta_t + \alpha_t^\theta \gamma^t \delta_t \nabla_\theta \log \pi_{\theta_t}(A_t | S_t)$

$w_{t+1} \leftarrow w_t + \alpha_t^w \delta_t \nabla_w \hat{v}(S_t, w_t)$

$S_t \leftarrow S_{t+1}$

until *converged*

1.11 The Ladder



1.12 What We Bought

QAC carried one critic (Q^{π_θ}) and paid for it in variance. This report carries one critic (V^{π_θ}) and less variance, from two moves that cost nothing in structural bias: subtracting the action-independent part of the signal (the baseline), and integrating out the next action analytically (Expected SARSA). Every estimator on the ladder is exactly unbiased; bias enters only in the last row of the table, and only through critic approximation $\hat{v} \neq V^{\pi_\theta}$, never through the estimator structure itself. That the bias column reads “none” until that final row is the whole point.

1.13 Scope

This report is one-step actor-critic. n -step actor-critic, eligibility traces, and GAE come next.

A Time-Marginal Form

The per-step estimators above inherit the γ^t -weighted time-marginal identity of Report 15, with the multiplier (Q^{π_θ} , or A^{π_θ}) carried through unchanged. For the advantage multiplier:

$$d^{\pi_\theta, \rho_0}(s) = (1 - \gamma) \sum_{t=0}^{\infty} \gamma^t \Pr_\theta(S_t = s).$$

$$\begin{aligned}
\nabla_{\theta} J(\theta) &= \frac{1}{1-\gamma} \sum_s d^{\pi_{\theta}, \rho_0}(s) \sum_a \pi_{\theta}(a \mid s) A^{\pi_{\theta}}(s, a) \nabla_{\theta} \log \pi_{\theta}(a \mid s) && \text{; advantage PG} \\
&= \frac{1}{1-\gamma} \sum_s \left[(1-\gamma) \sum_{t=0}^{\infty} \gamma^t \Pr_{\theta}(S_t = s) \right] \\
&\quad \cdot \sum_a \pi_{\theta}(a \mid s) A^{\pi_{\theta}}(s, a) \nabla_{\theta} \log \pi_{\theta}(a \mid s) && \text{; substitute } d^{\pi_{\theta}, \rho_0} \\
&= \sum_{t=0}^{\infty} \gamma^t \sum_s \Pr_{\theta}(S_t = s) \sum_a \pi_{\theta}(a \mid s) A^{\pi_{\theta}}(s, a) \nabla_{\theta} \log \pi_{\theta}(a \mid s) && \text{; cancel and rearrange} \\
&= \sum_{t=0}^{\infty} \gamma^t \mathbb{E}_{S_t \sim \Pr_{\theta}} \left[\mathbb{E}_{A_t \sim \pi_{\theta}(\cdot \mid S_t)} \left[A^{\pi_{\theta}}(S_t, A_t) \nabla_{\theta} \log \pi_{\theta}(A_t \mid S_t) \right] \right], && \text{; time-}t \text{ marginals}
\end{aligned}$$

so

$$\hat{g}_T^{A, \text{oracle}}(\theta) = \sum_{t=0}^{T-1} \gamma^t A^{\pi_{\theta}}(S_t, A_t) \nabla_{\theta} \log \pi_{\theta}(A_t \mid S_t).$$

B Estimator Naming Ladder

The four multipliers of the score, in order of refinement:

$$\begin{aligned}
\hat{g}_t^{Q, \text{oracle}} &= \gamma^t Q^{\pi_{\theta}}(S_t, A_t) \nabla_{\theta} \log \pi_{\theta}(A_t \mid S_t) \\
\hat{g}_t^{A, \text{oracle}} &= \gamma^t A^{\pi_{\theta}}(S_t, A_t) \nabla_{\theta} \log \pi_{\theta}(A_t \mid S_t) \\
\hat{g}_t^{\hat{A}} &= \gamma^t [\hat{q}(S_t, A_t, w_t^q) - \hat{v}(S_t, w_t^v)] \nabla_{\theta} \log \pi_{\theta}(A_t \mid S_t) \\
\hat{g}_t^{\delta} &= \gamma^t \delta_t \nabla_{\theta} \log \pi_{\theta}(A_t \mid S_t)
\end{aligned}$$