

Derivation of the Policy Gradient Formula

August 11, 2026

1 Setup

Consider a finite-horizon Markov decision process with horizon T . Let $\mathcal{S}$ and $\mathcal{A}$ denote the state and action spaces. For simplicity, all probability distributions below are written using densities or probability mass functions; the same derivation can be formulated with probability kernels and Radon–Nikodym derivatives.

Let

$$\tau = (S_0, A_0, S_1, A_1, \dots, S_{T-1}, A_{T-1}, S_T)$$

be the random trajectory. Its law is determined by

$$\begin{aligned} S_0 &\sim \rho_0, \\ A_t \mid S_t = s &\sim \pi_\theta(\cdot \mid s), \\ S_{t+1} \mid (S_t = s, A_t = a) &\sim P(\cdot \mid s, a), \end{aligned}$$

where π_θ is the parameterized policy. We assume that the initial-state distribution ρ_0 and the transition kernel P do not depend on θ .

Let the one-step reward be $R_t = r(S_t, A_t)$ and define the discounted return of a trajectory by

$$G(\tau) = \sum_{t=0}^{T-1} \gamma^t R_t, \quad 0 \leq \gamma \leq 1.$$

The objective is

$$J(\theta) = \mathbb{E}_{\tau \sim p_\theta}[G(\tau)],$$

where p_θ is the probability law of the entire trajectory induced jointly by the policy and the environment. Thus, the expectation in $J(\theta)$ is over all random states and actions in the trajectory.

The trajectory density factors as

$$p_\theta(\tau) = \rho_0(s_0) \prod_{t=0}^{T-1} \pi_\theta(a_t \mid s_t) P(s_{t+1} \mid s_t, a_t). \quad (1)$$

2 Likelihood-Ratio Representation

Writing the expectation as an integral gives

$$J(\theta) = \int G(\tau) p_\theta(\tau) \, d\tau.$$

Assume that differentiation may be interchanged with integration. Then

$$\begin{aligned}\nabla_{\theta} J(\theta) &= \int G(\tau) \nabla_{\theta} p_{\theta}(\tau) d\tau \\ &= \int G(\tau) p_{\theta}(\tau) \nabla_{\theta} \log p_{\theta}(\tau) d\tau \\ &= \mathbb{E}_{\tau \sim p_{\theta}} [G(\tau) \nabla_{\theta} \log p_{\theta}(\tau)].\end{aligned}$$

The second equality uses the score-function, or likelihood-ratio, identity

$$\nabla_{\theta} p_{\theta}(\tau) = p_{\theta}(\tau) \nabla_{\theta} \log p_{\theta}(\tau),$$

valid wherever $p_{\theta}(\tau) > 0$.

Taking logarithms in (1),

$$\begin{aligned}\log p_{\theta}(\tau) &= \log \rho_0(S_0) \\ &\quad + \sum_{t=0}^{T-1} (\log \pi_{\theta}(A_t | S_t) + \log P(S_{t+1} | S_t, A_t)).\end{aligned}$$

Since neither ρ_0 nor P depends on θ ,

$$\nabla_{\theta} \log p_{\theta}(\tau) = \sum_{t=0}^{T-1} \nabla_{\theta} \log \pi_{\theta}(A_t | S_t). \quad (2)$$

Consequently,

$$\begin{aligned}\nabla_{\theta} J(\theta) &= \mathbb{E}_{\tau \sim p_{\theta}} \left[\left(\sum_{k=0}^{T-1} \gamma^k R_k \right) \left(\sum_{t=0}^{T-1} \nabla_{\theta} \log \pi_{\theta}(A_t | S_t) \right) \right] \\ &= \sum_{t=0}^{T-1} \mathbb{E}_{\tau \sim p_{\theta}} \left[\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) \sum_{k=0}^{T-1} \gamma^k R_k \right].\end{aligned} \quad (3)$$

3 Removal of Past Rewards

For a fixed time t , define the information available immediately before A_t is sampled by

$$\mathcal{F}_t = \sigma(S_0, A_0, \dots, A_{t-1}, S_t).$$

The random variable

$$\sum_{k=0}^{t-1} \gamma^k R_k$$

is $\mathcal{F}_t$ -measurable. We now show that its contribution to (3) is zero.

Conditionally on $\mathcal{F}_t$, the state S_t is known and the only new randomness in the score term comes from

$$A_t \sim \pi_{\theta}(\cdot | S_t).$$

Hence

$$\begin{aligned}\mathbb{E}_{\tau \sim p_{\theta}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) | \mathcal{F}_t] &= \int \pi_{\theta}(a | S_t) \nabla_{\theta} \log \pi_{\theta}(a | S_t) da \\ &= \int \nabla_{\theta} \pi_{\theta}(a | S_t) da \\ &= \nabla_{\theta} \int \pi_{\theta}(a | S_t) da = 0.\end{aligned}$$

Therefore, by the tower property of conditional expectation,

$$\begin{aligned} & \mathbb{E}_{\tau \sim p_\theta} \left[\nabla_\theta \log \pi_\theta(A_t \mid S_t) \sum_{k=0}^{t-1} \gamma^k R_k \right] \\ &= \mathbb{E}_{\tau \sim p_\theta} \left[\left(\sum_{k=0}^{t-1} \gamma^k R_k \right) \mathbb{E}_{\tau \sim p_\theta} [\nabla_\theta \log \pi_\theta(A_t \mid S_t) \mid \mathcal{F}_t] \right] = 0. \end{aligned}$$

Thus only rewards at time t and later remain:

$$\nabla_\theta J(\theta) = \sum_{t=0}^{T-1} \mathbb{E}_{\tau \sim p_\theta} \left[\nabla_\theta \log \pi_\theta(A_t \mid S_t) \sum_{k=t}^{T-1} \gamma^k R_k \right]. \quad (4)$$

Define the return measured from time t by

$$G_t = \sum_{\ell=0}^{T-1-t} \gamma^\ell R_{t+\ell}.$$

Since

$$\sum_{k=t}^{T-1} \gamma^k R_k = \gamma^t G_t,$$

we obtain

$$\nabla_\theta J(\theta) = \sum_{t=0}^{T-1} \gamma^t \mathbb{E}_{\tau \sim p_\theta} [\nabla_\theta \log \pi_\theta(A_t \mid S_t) G_t]. \quad (5)$$

This is the REINFORCE form of the policy gradient.

4 Conditional Expectation and the Action-Value Function

Define the finite-horizon action-value function at time t by

$$Q_t^{\pi_\theta}(s, a) = \mathbb{E}_{\pi_\theta} [G_t \mid S_t = s, A_t = a]. \quad (6)$$

The conditional expectation in (6) is taken over the future randomness after the state-action pair $(S_t, A_t) = (s, a)$ is fixed. In particular, it averages over future state transitions and over future actions sampled from π_θ .

Since

$$\nabla_\theta \log \pi_\theta(A_t \mid S_t)$$

is measurable with respect to $\sigma(S_t, A_t)$, the tower property gives

$$\begin{aligned} & \mathbb{E}_{\tau \sim p_\theta} [\nabla_\theta \log \pi_\theta(A_t \mid S_t) G_t] \\ &= \mathbb{E}_{\tau \sim p_\theta} [\nabla_\theta \log \pi_\theta(A_t \mid S_t) \mathbb{E}_{\tau \sim p_\theta} [G_t \mid S_t, A_t]] \\ &= \mathbb{E}_{(S_t, A_t) \sim p_{\theta,t}} [\nabla_\theta \log \pi_\theta(A_t \mid S_t) Q_t^{\pi_\theta}(S_t, A_t)], \end{aligned}$$

where $p_{\theta,t}$ denotes the marginal law of (S_t, A_t) under the trajectory distribution p_θ .

Let $d_t^{\pi_\theta}$ denote the marginal law of S_t . Then $p_{\theta,t}$ has the factorization

$$p_{\theta,t}(ds, da) = d_t^{\pi_\theta}(ds) \pi_\theta(da \mid s),$$

and, for every integrable function f ,

$$\begin{aligned} & \mathbb{E}_{(S_t, A_t) \sim p_{\theta,t}} [f(S_t, A_t)] \\ &= \int_{\mathcal{S}} d_t^{\pi_\theta}(ds) \int_{\mathcal{A}} \pi_\theta(da \mid s) f(s, a). \end{aligned}$$

Thus the policy gradient becomes

$$\nabla_{\theta} J(\theta) = \sum_{t=0}^{T-1} \gamma^t \mathbb{E}_{(S_t, A_t) \sim p_{\theta, t}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) Q_t^{\pi_{\theta}}(S_t, A_t)]. \quad (7)$$

The expectation in (7) is over two levels of randomness: first S_t is distributed according to the state distribution generated by the policy and environment up to time t ; then, conditional on $S_t = s$, the action is sampled from $\pi_{\theta}(\cdot | s)$.

5 State-Dependent Baselines and Advantage

Let $b_t : \mathcal{S} \rightarrow \mathbb{R}$ be any integrable function that does not depend on the sampled action. We have

$$\begin{aligned} & \mathbb{E}_{(S_t, A_t) \sim p_{\theta, t}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) b_t(S_t)] \\ &= \mathbb{E}_{S_t \sim d_t^{\pi_{\theta}}} [b_t(S_t) \mathbb{E}_{A_t \sim \pi_{\theta}(\cdot | S_t)} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t)]] = 0. \end{aligned}$$

Therefore subtracting a state-dependent baseline does not change the expected policy gradient:

$$\begin{aligned} & \mathbb{E}_{(S_t, A_t) \sim p_{\theta, t}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) Q_t^{\pi_{\theta}}(S_t, A_t)] \\ &= \mathbb{E}_{(S_t, A_t) \sim p_{\theta, t}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) (Q_t^{\pi_{\theta}}(S_t, A_t) - b_t(S_t))]. \end{aligned}$$

Define the state-value function by

$$V_t^{\pi_{\theta}}(s) = \mathbb{E}_{A \sim \pi_{\theta}(\cdot | s)} [Q_t^{\pi_{\theta}}(s, A)]$$

and the advantage function by

$$A_t^{\pi_{\theta}}(s, a) = Q_t^{\pi_{\theta}}(s, a) - V_t^{\pi_{\theta}}(s).$$

Choosing $b_t = V_t^{\pi_{\theta}}$ in (7) yields

$$\nabla_{\theta} J(\theta) = \sum_{t=0}^{T-1} \gamma^t \mathbb{E}_{(S_t, A_t) \sim p_{\theta, t}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) A_t^{\pi_{\theta}}(S_t, A_t)]. \quad (8)$$

The baseline changes the variance of a Monte Carlo gradient estimator but not its expectation. The value function is a particularly natural baseline because the advantage measures the value of the sampled action relative to the policy's average action value at the same state.

6 Infinite-Horizon Discounted Form

Now suppose $T = \infty$ and $0 \leq \gamma < 1$. Under standard integrability and regularity assumptions allowing the preceding exchanges of infinite sums, expectations, and derivatives, define the normalized discounted state-occupancy measure

$$d_{\pi_{\theta}}^{\gamma} = (1 - \gamma) \sum_{t=0}^{\infty} \gamma^t \text{Law}_{\pi_{\theta}}(S_t). \quad (9)$$

This is a probability measure because each $\text{Law}_{\pi_\theta}(S_t)$ has total mass one and

$$(1 - \gamma) \sum_{t=0}^{\infty} \gamma^t = 1.$$

If $S \sim d_{\pi_\theta}'$ and, conditionally on $S = s$, $A \sim \pi_\theta(\cdot | s)$, then

$$\boxed{\nabla_\theta J(\theta) = \frac{1}{1 - \gamma} \mathbb{E}_{\substack{S \sim d_{\pi_\theta}' \\ A \sim \pi_\theta(\cdot | S)}} [\nabla_\theta \log \pi_\theta(A | S) A^{\pi_\theta}(S, A)]}. \quad (10)$$

Equation (10) makes precise the common shorthand

$$\nabla_\theta J(\theta) \propto \mathbb{E}_{s, a \sim \pi_\theta} [\nabla_\theta \log \pi_\theta(a | s) A^{\pi_\theta}(s, a)].$$

The notation $s, a \sim \pi_\theta$ suppresses two distinct distributions. The state is distributed according to the discounted occupancy measure induced by repeated interaction of π_θ with the environment, while the action is then sampled conditionally from $\pi_\theta(\cdot | s)$.