

Natural Policy Gradient

Reinforcement Learning

Sai Sampath Kedari

sampath@umich.edu

Contents

1	Local Models, Trust Regions, and Vanilla Policy Gradient	2
1.1	Trust-Region Viewpoint	2
1.2	Euclidean Trust-Region Model	2
1.3	Lagrangian Solution	3
1.4	Vanilla Policy-Gradient Update	4
1.5	Why Euclidean Geometry Is Insufficient	5
2	Policy-Space Geometry and the KL Constraint	5
2.1	KL-Constrained Policy Update	5
2.2	First-Order Approximation of the Objective	5
2.3	Second-Order Approximation of the KL Constraint	6
2.4	Local Natural Policy-Gradient Problem	8
3	Solving the Natural Policy-Gradient Problem	9
3.1	Lagrangian Solution	9
3.2	Natural Policy-Gradient Updates	10
4	Practical Computation of the NPG Update	10
4.1	Damped Natural-Gradient Direction	11
4.2	Fisher-Vector Products and Conjugate Gradient	11
4.3	Choosing the Step Length	11
4.4	Practical NPG Algorithm	12

1 Local Models, Trust Regions, and Vanilla Policy Gradient

1.1 Trust-Region Viewpoint

Suppose the objective is

$$J(\theta).$$

Directly solving

$$\theta^* = \arg \max_{\theta} J(\theta)$$

is generally difficult because J may be highly nonlinear and only accessible through interaction with the environment.

Trust-region optimization instead proceeds locally. At the current parameters θ_{old} , we construct a simpler surrogate model of the objective using local information available at θ_{old} .

Define the candidate displacement

$$\Delta\theta := \theta - \theta_{\text{old}}, \quad \theta = \theta_{\text{old}} + \Delta\theta.$$

A local model

$$m_{\text{old}}(\Delta\theta) \approx J(\theta_{\text{old}} + \Delta\theta)$$

is expected to be accurate near $\Delta\theta = 0$, but there is no reason to trust the approximation arbitrarily far from the expansion point.

We therefore optimize the surrogate only inside a neighborhood around the current parameters:

$$\boxed{\begin{array}{ll} \Delta\theta^* = \arg \max_{\Delta\theta} & m_{\text{old}}(\Delta\theta) \\ \text{subject to} & \text{size}(\Delta\theta) \leq \delta. \end{array}}$$

The constraint defines the *trust region*: the region in which the local surrogate is trusted as an approximation to the true objective.

The particular choice of $\text{size}(\Delta\theta)$ determines the geometry of the trust region.

1.2 Euclidean Trust-Region Model

From the previous report, the policy objective is

$$J(\theta) = \mathbb{E}_{S_0 \sim \rho_0} [v^{\pi_\theta}(S_0)],$$

with policy gradient

$$\nabla_{\theta} J(\theta) = \mathbb{E}_{\substack{S_t \sim d^{\pi_{\theta}, \rho_0} \\ A_t \sim \pi_{\theta}(\cdot | S_t)}} [\nabla_{\theta} \log \pi_{\theta}(A_t | S_t) A^{\pi_{\theta}}(S_t, A_t)].$$

At θ_{old} , define

$$g_{\text{old}} := \nabla_{\theta} J(\theta)|_{\theta=\theta_{\text{old}}}.$$

The first-order Taylor expansion is

$$J(\theta_{\text{old}} + \Delta\theta) = J(\theta_{\text{old}}) + g_{\text{old}}^{\top} \Delta\theta + o(\|\Delta\theta\|_2).$$

Dropping the higher-order term gives the first-order surrogate

$$m_{\text{old}}(\Delta\theta) = J(\theta_{\text{old}}) + g_{\text{old}}^\top \Delta\theta.$$

Since $J(\theta_{\text{old}})$ is constant with respect to $\Delta\theta$, maximizing the surrogate is equivalent to maximizing

$$g_{\text{old}}^\top \Delta\theta.$$

Using the Euclidean norm to define the trust region gives

$$\begin{aligned} \Delta\theta^* &= \arg \max_{\Delta\theta} g_{\text{old}}^\top \Delta\theta \\ \text{subject to } & \frac{1}{2} \Delta\theta^\top \Delta\theta \leq \delta. \end{aligned}$$

For $g_{\text{old}} \neq 0$, the optimum lies on the boundary,

$$\frac{1}{2} \Delta\theta^{*\top} \Delta\theta^* = \delta.$$

1.3 Lagrangian Solution

Define

$$\mathcal{L}(\Delta\theta, \eta) = g_{\text{old}}^\top \Delta\theta - \eta \left(\frac{1}{2} \Delta\theta^\top \Delta\theta - \delta \right), \quad \eta \geq 0.$$

Taking the total differential,

$$\begin{aligned} d\mathcal{L} &= d(g_{\text{old}}^\top \Delta\theta) - d \left[\eta \left(\frac{1}{2} \Delta\theta^\top \Delta\theta - \delta \right) \right] \\ &= g_{\text{old}}^\top d(\Delta\theta) - \left(\frac{1}{2} \Delta\theta^\top \Delta\theta - \delta \right) d\eta - \frac{\eta}{2} d(\Delta\theta^\top \Delta\theta) \\ &= g_{\text{old}}^\top d(\Delta\theta) - \left(\frac{1}{2} \Delta\theta^\top \Delta\theta - \delta \right) d\eta \\ &\quad - \frac{\eta}{2} \left[d(\Delta\theta)^\top \Delta\theta + \Delta\theta^\top d(\Delta\theta) \right] \\ &= \left(g_{\text{old}}^\top - \eta \Delta\theta^\top \right) d(\Delta\theta) + \left(\delta - \frac{1}{2} \Delta\theta^\top \Delta\theta \right) d\eta, \end{aligned}$$

where

$$d(\Delta\theta)^\top \Delta\theta = \Delta\theta^\top d(\Delta\theta),$$

since both sides are scalars.

Comparing with

$$d\mathcal{L} = \nabla_{\Delta\theta} \mathcal{L}^\top d(\Delta\theta) + \frac{\partial \mathcal{L}}{\partial \eta} d\eta,$$

gives

$$\nabla_{\Delta\theta} \mathcal{L} = g_{\text{old}} - \eta \Delta\theta, \quad \frac{\partial \mathcal{L}}{\partial \eta} = \delta - \frac{1}{2} \Delta\theta^\top \Delta\theta.$$

Setting both equal to zero,

$$\Delta\theta^\star = \frac{1}{\eta}g_{\text{old}}, \quad \frac{1}{2}\Delta\theta^{\star\top}\Delta\theta^\star = \delta.$$

Therefore,

$$\frac{1}{2\eta^2}g_{\text{old}}^\top g_{\text{old}} = \delta,$$

and hence

$$\eta^\star = \sqrt{\frac{g_{\text{old}}^\top g_{\text{old}}}{2\delta}}.$$

Substituting back,

$$\Delta\theta^\star = \sqrt{\frac{2\delta}{g_{\text{old}}^\top g_{\text{old}}}}g_{\text{old}} = \sqrt{2\delta} \frac{g_{\text{old}}}{\|g_{\text{old}}\|_2}.$$

Thus,

$$\frac{g_{\text{old}}}{\|g_{\text{old}}\|_2}$$

is the normalized Euclidean steepest-ascent direction, while

$$g_{\text{old}}$$

is the corresponding unnormalized direction.

1.4 Vanilla Policy-Gradient Update

Rather than fixing the Euclidean trust-region radius explicitly, ordinary gradient ascent controls the step magnitude with a learning rate $\alpha > 0$:

$$\Delta\theta = \alpha g_{\text{old}}.$$

Therefore,

$$\theta_{\text{new}} = \theta_{\text{old}} + \alpha g_{\text{old}}.$$

Using the sample policy-gradient estimate

$$\hat{g}_{\text{old}} = \hat{\mathbb{E}}_t \left[\nabla_\theta \log \pi_\theta(A_t \mid S_t) \big|_{\theta=\theta_{\text{old}}} \hat{A}_t^{\text{GAE}(\gamma, \lambda)} \right],$$

the practical vanilla policy-gradient update is

$$\theta_{\text{new}} = \theta_{\text{old}} + \alpha \hat{g}_{\text{old}}.$$

1.5 Why Euclidean Geometry Is Insufficient

The trust-region idea itself is not the problem. The questionable choice is the Euclidean measure

$$\frac{1}{2}\Delta\theta^\top\Delta\theta.$$

It controls how far the *parameters* move, but the object being optimized is the policy distribution

$$\pi_\theta(\cdot | s).$$

Equal Euclidean parameter displacements can induce very different changes in the corresponding action distributions. Therefore, Euclidean parameter distance does not provide a reliable measure of how much the policy itself has changed.

Natural policy gradient keeps the same local-model and trust-region viewpoint, but replaces Euclidean parameter-space geometry with a geometry defined directly by changes in the policy distribution.

2 Policy-Space Geometry and the KL Constraint

2.1 KL-Constrained Policy Update

At the current parameters θ_{old} , consider

$$\theta = \theta_{\text{old}} + \Delta\theta.$$

Define the average KL divergence

$$\bar{D}_{\text{KL}}(\pi_{\text{old}}\|\pi_\theta) := \mathbb{E}_{S \sim d^{\pi_{\text{old}}, \rho_0}} [D_{\text{KL}}(\pi_{\text{old}}(\cdot | S)\|\pi_\theta(\cdot | S))].$$

For a fixed state S ,

$$D_{\text{KL}}(\pi_{\text{old}}(\cdot | S)\|\pi_\theta(\cdot | S)) = \mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\log \frac{\pi_{\text{old}}(A | S)}{\pi_\theta(A | S)} \right].$$

The policy-space trust-region problem is

$$\begin{aligned} \Delta\theta^* &= \arg \max_{\Delta\theta} J(\theta_{\text{old}} + \Delta\theta) \\ \text{subject to} \quad &\bar{D}_{\text{KL}}(\pi_{\text{old}}\|\pi_{\theta_{\text{old}} + \Delta\theta}) \leq \delta. \end{aligned}$$

2.2 First-Order Approximation of the Objective

Using the first-order model already derived,

$$J(\theta_{\text{old}} + \Delta\theta) \approx J(\theta_{\text{old}}) + g_{\text{old}}^\top \Delta\theta.$$

Since $J(\theta_{\text{old}})$ is constant,

$$\begin{aligned} \Delta\theta^* &\approx \arg \max_{\Delta\theta} g_{\text{old}}^\top \Delta\theta \\ \text{subject to} \quad &\bar{D}_{\text{KL}}(\pi_{\text{old}}\|\pi_{\theta_{\text{old}} + \Delta\theta}) \leq \delta. \end{aligned}$$

2.3 Second-Order Approximation of the KL Constraint

For compactness, define

$$\bar{D}_{\text{KL}}(\theta) := \bar{D}_{\text{KL}}(\pi_{\text{old}} \parallel \pi_{\theta}).$$

Expanding around θ_{old} ,

$$\begin{aligned} \bar{D}_{\text{KL}}(\theta_{\text{old}} + \Delta\theta) &\approx \underbrace{\bar{D}_{\text{KL}}(\theta_{\text{old}})}_{\text{0th-order term}} \\ &\quad + \underbrace{\nabla_{\theta} \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}}^{\top} \Delta\theta}_{\text{1st-order term}} \\ &\quad + \underbrace{\frac{1}{2} \Delta\theta^{\top} \nabla_{\theta}^2 \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}} \Delta\theta}_{\text{2nd-order term}}. \end{aligned}$$

0th-order term. At the expansion point, $\theta = \theta_{\text{old}}$, so the candidate policy π_{θ} becomes the old policy π_{old} . Therefore,

$$\begin{aligned} \bar{D}_{\text{KL}}(\theta_{\text{old}}) &= \bar{D}_{\text{KL}}(\pi_{\text{old}} \parallel \pi_{\text{old}}) \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} [D_{\text{KL}}(\pi_{\text{old}}(\cdot | S) \parallel \pi_{\text{old}}(\cdot | S))] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} [0] \\ &= 0. \end{aligned}$$

Hence, the zeroth-order term in the Taylor expansion is

$$\boxed{\bar{D}_{\text{KL}}(\theta_{\text{old}}) = 0.}$$

1st-order term. The gradient of the average KL divergence at the expansion point is

$$\begin{aligned} \nabla_{\theta} \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}} &= \nabla_{\theta} \bar{D}_{\text{KL}}(\pi_{\text{old}} \parallel \pi_{\theta}) \big|_{\theta=\theta_{\text{old}}} \\ &= \nabla_{\theta} \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} [D_{\text{KL}}(\pi_{\text{old}}(\cdot | S) \parallel \pi_{\theta}(\cdot | S))] \big|_{\theta=\theta_{\text{old}}} \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\nabla_{\theta} D_{\text{KL}}(\pi_{\text{old}}(\cdot | S) \parallel \pi_{\theta}(\cdot | S)) \big|_{\theta=\theta_{\text{old}}} \right] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\nabla_{\theta} \mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\log \frac{\pi_{\text{old}}(A | S)}{\pi_{\theta}(A | S)} \right] \bigg|_{\theta=\theta_{\text{old}}} \right] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\nabla_{\theta} (\log \pi_{\text{old}}(A | S) - \log \pi_{\theta}(A | S)) \big|_{\theta=\theta_{\text{old}}} \right] \right] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\underbrace{\nabla_{\theta} \log \pi_{\text{old}}(A | S)}_0 - \nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \right] \right] \\ &= -\mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \right] \right]. \end{aligned}$$

By the zero-mean score-function identity,

$$\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \right] = 0.$$

Therefore,

$$\begin{aligned}\nabla_{\theta} \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}} &= -\mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} [0] \\ &= 0.\end{aligned}$$

Hence,

$$\boxed{\nabla_{\theta} \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}} = 0.}$$

Consequently, the complete first-order Taylor term vanishes:

$$\boxed{\nabla_{\theta} \bar{D}_{\text{KL}}(\theta)^{\top} \big|_{\theta=\theta_{\text{old}}} \Delta\theta = 0.}$$

The zero-mean score identity also has a maximum-likelihood interpretation. For each fixed state S ,

$$\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} [\log \pi_{\theta}(A | S)]$$

is the population expected log-likelihood under actions generated by π_{old} . It is maximized when the candidate distribution matches the data-generating distribution, which occurs at $\theta = \theta_{\text{old}}$. Therefore, its gradient at θ_{old} is zero.

2nd-order term. The Hessian of the average KL divergence at the expansion point is

$$\begin{aligned}\nabla_{\theta}^2 \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}} &= \nabla_{\theta}^2 \bar{D}_{\text{KL}}(\pi_{\text{old}} \| \pi_{\theta}) \big|_{\theta=\theta_{\text{old}}} \\ &= \nabla_{\theta}^2 \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} [D_{\text{KL}}(\pi_{\text{old}}(\cdot | S) \| \pi_{\theta}(\cdot | S))] \big|_{\theta=\theta_{\text{old}}} \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\nabla_{\theta}^2 D_{\text{KL}}(\pi_{\text{old}}(\cdot | S) \| \pi_{\theta}(\cdot | S)) \big|_{\theta=\theta_{\text{old}}} \right] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\nabla_{\theta}^2 \mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\log \frac{\pi_{\text{old}}(A | S)}{\pi_{\theta}(A | S)} \right] \bigg|_{\theta=\theta_{\text{old}}} \right] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\nabla_{\theta}^2 (\log \pi_{\text{old}}(A | S) - \log \pi_{\theta}(A | S)) \big|_{\theta=\theta_{\text{old}}} \right] \right] \\ &= \mathbb{E}_{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0}} \left[\mathbb{E}_{A \sim \pi_{\text{old}}(\cdot | S)} \left[\underbrace{\nabla_{\theta}^2 \log \pi_{\text{old}}(A | S)}_0 - \nabla_{\theta}^2 \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \right] \right] \\ &= -\mathbb{E}_{\substack{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0} \\ A \sim \pi_{\text{old}}(\cdot | S)}} \left[\nabla_{\theta}^2 \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \right].\end{aligned}$$

By the Fisher information identity,

$$\begin{aligned}& -\mathbb{E}_{\substack{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0} \\ A \sim \pi_{\text{old}}(\cdot | S)}} \left[\nabla_{\theta}^2 \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \right] \\ &= \mathbb{E}_{\substack{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0} \\ A \sim \pi_{\text{old}}(\cdot | S)}} \left[\nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}}^{\top} \right].\end{aligned}$$

Therefore, define the Fisher information matrix at the old policy as

$$\boxed{\begin{aligned}F_{\text{old}} &:= \nabla_{\theta}^2 \bar{D}_{\text{KL}}(\theta) \big|_{\theta=\theta_{\text{old}}} \\ &= \mathbb{E}_{\substack{S \sim d^{\pi_{\theta_{\text{old}}}, \rho_0} \\ A \sim \pi_{\text{old}}(\cdot | S)}} \left[\nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}} \nabla_{\theta} \log \pi_{\theta}(A | S) \big|_{\theta=\theta_{\text{old}}}^{\top} \right].\end{aligned}}$$

Substituting the 0th, 1st, and 2nd order terms into the Taylor expansion gives the following.

$$\begin{aligned}\bar{D}_{\text{KL}}(\pi_{\text{old}} \parallel \pi_{\theta_{\text{old}} + \Delta\theta}) &\approx \underbrace{\bar{D}_{\text{KL}}(\theta_{\text{old}})}_0 + \underbrace{\nabla_{\theta} \bar{D}_{\text{KL}}(\theta)|_{\theta=\theta_{\text{old}}}}_0^{\top} \Delta\theta + \frac{1}{2} \Delta\theta^{\top} \underbrace{\nabla_{\theta}^2 \bar{D}_{\text{KL}}(\theta)|_{\theta=\theta_{\text{old}}}}_{F_{\text{old}}} \Delta\theta \\ &= \frac{1}{2} \Delta\theta^{\top} F_{\text{old}} \Delta\theta.\end{aligned}$$

Hence, the second-order local approximation of the average KL divergence is

$$\boxed{\bar{D}_{\text{KL}}(\pi_{\text{old}} \parallel \pi_{\theta_{\text{old}} + \Delta\theta}) \approx \frac{1}{2} \Delta\theta^{\top} F_{\text{old}} \Delta\theta.}$$

2.4 Local Natural Policy-Gradient Problem

The problem we would like to solve is

$$\boxed{\theta^* = \arg \max_{\theta} J(\theta).}$$

$\Downarrow$ construct a local trust region in policy-distribution space

$$\boxed{\begin{aligned} \Delta\theta^* &= \arg \max_{\Delta\theta} J(\theta_{\text{old}} + \Delta\theta) \\ \text{subject to } &\bar{D}_{\text{KL}}(\pi_{\text{old}} \parallel \pi_{\theta_{\text{old}} + \Delta\theta}) \leq \delta. \end{aligned}}$$

$\Downarrow$ first-order model of J and second-order model of $\bar{D}_{\text{KL}}$

$$\boxed{\begin{aligned} \Delta\theta^* &= \arg \max_{\Delta\theta} J(\theta_{\text{old}}) + g_{\text{old}}^{\top} \Delta\theta \\ \text{subject to } &\frac{1}{2} \Delta\theta^{\top} F_{\text{old}} \Delta\theta \leq \delta. \end{aligned}}$$

$\Downarrow$ $J(\theta_{\text{old}})$ is constant with respect to $\Delta\theta$

$$\boxed{\begin{aligned} \Delta\theta^* &= \arg \max_{\Delta\theta} g_{\text{old}}^{\top} \Delta\theta \\ \text{subject to } &\frac{1}{2} \Delta\theta^{\top} F_{\text{old}} \Delta\theta \leq \delta. \end{aligned}}$$

This local problem has the same structure as the Euclidean trust-region problem of Section 1. The matrix defining the local trust-region geometry changes from the identity matrix to the Fisher information matrix:

$$\boxed{I \longrightarrow F_{\text{old}}.}$$

3 Solving the Natural Policy-Gradient Problem

3.1 Lagrangian Solution

For $g_{\text{old}} \neq 0$, the constraint is active:

$$\frac{1}{2}\Delta\theta^{\star\top}F_{\text{old}}\Delta\theta^{\star} = \delta.$$

Define

$$\mathcal{L}(\Delta\theta, \eta) = g_{\text{old}}^{\top}\Delta\theta - \eta \left(\frac{1}{2}\Delta\theta^{\top}F_{\text{old}}\Delta\theta - \delta \right), \quad \eta \geq 0.$$

Taking the total differential,

$$\begin{aligned} d\mathcal{L} &= d(g_{\text{old}}^{\top}\Delta\theta) - d \left[\eta \left(\frac{1}{2}\Delta\theta^{\top}F_{\text{old}}\Delta\theta - \delta \right) \right] \\ &= g_{\text{old}}^{\top}d(\Delta\theta) - \left(\frac{1}{2}\Delta\theta^{\top}F_{\text{old}}\Delta\theta - \delta \right) d\eta - \frac{\eta}{2}d \left(\Delta\theta^{\top}F_{\text{old}}\Delta\theta \right) \\ &= g_{\text{old}}^{\top}d(\Delta\theta) - \left(\frac{1}{2}\Delta\theta^{\top}F_{\text{old}}\Delta\theta - \delta \right) d\eta \\ &\quad - \frac{\eta}{2} \left[d(\Delta\theta)^{\top}F_{\text{old}}\Delta\theta + \Delta\theta^{\top}F_{\text{old}}d(\Delta\theta) \right] \\ &= \left(g_{\text{old}}^{\top} - \eta\Delta\theta^{\top}F_{\text{old}} \right) d(\Delta\theta) + \left(\delta - \frac{1}{2}\Delta\theta^{\top}F_{\text{old}}\Delta\theta \right) d\eta, \end{aligned}$$

where $F_{\text{old}}^{\top} = F_{\text{old}}$.

Comparing with

$$d\mathcal{L} = \nabla_{\Delta\theta}\mathcal{L}^{\top}d(\Delta\theta) + \frac{\partial\mathcal{L}}{\partial\eta}d\eta,$$

gives

$$\nabla_{\Delta\theta}\mathcal{L} = g_{\text{old}} - \eta F_{\text{old}}\Delta\theta,$$

and

$$\frac{\partial\mathcal{L}}{\partial\eta} = \delta - \frac{1}{2}\Delta\theta^{\top}F_{\text{old}}\Delta\theta.$$

Setting both equal to zero,

$$g_{\text{old}} - \eta F_{\text{old}}\Delta\theta^{\star} = 0,$$

so

$$\Delta\theta^{\star} = \frac{1}{\eta}F_{\text{old}}^{-1}g_{\text{old}}.$$

Hence,

$$\tilde{g}_{\text{old}} := F_{\text{old}}^{-1}g_{\text{old}}$$

is the unnormalized natural policy-gradient direction.

Substituting into the active constraint,

$$\frac{1}{2\eta^2}g_{\text{old}}^{\top}F_{\text{old}}^{-1}g_{\text{old}} = \delta,$$

which gives

$$\eta^* = \sqrt{\frac{g_{\text{old}}^\top F_{\text{old}}^{-1} g_{\text{old}}}{2\delta}}.$$

Therefore,

$$\Delta\theta^* = \sqrt{\frac{2\delta}{g_{\text{old}}^\top F_{\text{old}}^{-1} g_{\text{old}}}} F_{\text{old}}^{-1} g_{\text{old}}.$$

This is the normalized natural policy-gradient step.

3.2 Natural Policy-Gradient Updates

With an arbitrary learning rate $\alpha > 0$,

$$\theta_{\text{new}} = \theta_{\text{old}} + \alpha F_{\text{old}}^{-1} g_{\text{old}}.$$

Using sample estimates,

$$\theta_{\text{new}} = \theta_{\text{old}} + \alpha \hat{F}_{\text{old}}^{-1} \hat{g}_{\text{old}}.$$

If the KL trust-region radius directly determines the step magnitude,

$$\theta_{\text{new}} = \theta_{\text{old}} + \sqrt{\frac{2\delta}{\hat{g}_{\text{old}}^\top \hat{F}_{\text{old}}^{-1} \hat{g}_{\text{old}}}} \hat{F}_{\text{old}}^{-1} \hat{g}_{\text{old}}.$$

4 Practical Computation of the NPG Update

Consider a fixed rollout batch generated by π_{old} :

$$\mathcal{B}_{\text{old}} = \left\{ S_i, A_i, \hat{A}_i^{\pi_{\text{old}}} \right\}_{i=1}^B.$$

The empirical policy-gradient estimate is

$$\hat{g}_{\text{old}} = \frac{1}{B} \sum_{i=1}^B \nabla_{\theta} \log \pi_{\theta}(A_i \mid S_i) \Big|_{\theta=\theta_{\text{old}}} \hat{A}_i^{\pi_{\text{old}}}.$$

The empirical average KL divergence is

$$\widehat{\bar{D}}_{\text{KL}}(\theta) := \frac{1}{B} \sum_{i=1}^B D_{\text{KL}}(\pi_{\text{old}}(\cdot \mid S_i) \parallel \pi_{\theta}(\cdot \mid S_i)).$$

Its Hessian at the old policy defines the empirical Fisher matrix:

$$\hat{F}_{\text{old}} = \nabla_{\theta}^2 \widehat{\bar{D}}_{\text{KL}}(\theta) \Big|_{\theta=\theta_{\text{old}}}.$$

4.1 Damped Natural-Gradient Direction

The empirical Fisher matrix may be singular or poorly conditioned. Define

$$\tilde{F}_{\text{old}} = \hat{F}_{\text{old}} + \lambda_{\text{damp}} I, \quad \lambda_{\text{damp}} > 0.$$

Rather than explicitly computing $\tilde{F}_{\text{old}}^{-1}$, solve

$$\tilde{F}_{\text{old}} x = \hat{g}_{\text{old}}.$$

The solution

$$x = \tilde{F}_{\text{old}}^{-1} \hat{g}_{\text{old}}$$

is the damped natural-gradient direction.

4.2 Fisher-Vector Products and Conjugate Gradient

Constructing $\hat{F}_{\text{old}}$ explicitly is impractical for a neural-network policy. Conjugate gradient avoids this by requiring only products of the form $\hat{F}_{\text{old}} v$.

The Fisher-vector product is

$$\hat{F}_{\text{old}} v = \nabla_{\theta} \left[\left(\nabla_{\theta} \widehat{D}_{\text{KL}}(\theta) \right)^{\top} v \right] \Big|_{\theta=\theta_{\text{old}}}.$$

Therefore, the damped matrix-vector product used by conjugate gradient is

$$\tilde{F}_{\text{old}} v = \hat{F}_{\text{old}} v + \lambda_{\text{damp}} v.$$

After K conjugate-gradient iterations,

$$x_K \approx \tilde{F}_{\text{old}}^{-1} \hat{g}_{\text{old}}.$$

No Fisher matrix is constructed or inverted explicitly.

4.3 Choosing the Step Length

For fixed-learning-rate NPG, the update is

$$\theta_{\text{new}} = \theta_{\text{old}} + \alpha x_K.$$

For KL-normalized NPG, write

$$\Delta \theta_{\text{NPG}} = \alpha_{\delta} x_K.$$

The original quadratic KL model is

$$\frac{1}{2} \Delta \theta_{\text{NPG}}^{\top} \hat{F}_{\text{old}} \Delta \theta_{\text{NPG}} \leq \delta.$$

Placing the proposed step on its boundary gives

$$\frac{1}{2}\alpha_\delta^2 x_K^\top \widehat{F}_{\text{old}} x_K = \delta.$$

Therefore,

$$\alpha_\delta = \sqrt{\frac{2\delta}{x_K^\top \widehat{F}_{\text{old}} x_K}}.$$

The normalized NPG update is consequently

$$\theta_{\text{new}} = \theta_{\text{old}} + \sqrt{\frac{2\delta}{x_K^\top \widehat{F}_{\text{old}} x_K}} x_K.$$

Here, damping regularizes the direction computation, while the final scaling uses $\widehat{F}_{\text{old}}$ because $\widehat{F}_{\text{old}}$ defines the original quadratic KL model.

4.4 Practical NPG Algorithm

The following algorithm computes the KL-normalized NPG update on a fixed old-policy rollout batch.

Algorithm 1 KL-Normalized Natural Policy-Gradient Update

Require: Fixed rollout batch $\mathcal{B}_{\text{old}}$

Require: Old parameters θ_{old}

Require: KL radius δ

Require: Damping coefficient λ_{damp}

Require: Number of conjugate-gradient iterations K

- 1: Compute the empirical policy gradient $\widehat{g}_{\text{old}}$
 - 2: Define $\text{FVP}(v) \leftarrow \widehat{F}_{\text{old}} v$
 - 3: Define $\text{DampedFVP}(v) \leftarrow \text{FVP}(v) + \lambda_{\text{damp}} v$
 - 4: $x_K \leftarrow \text{CONJUGATEGRADIENT}(\text{DampedFVP}, \widehat{g}_{\text{old}}, K)$
 - 5: $q_K \leftarrow x_K^\top \text{FVP}(x_K)$
 - 6: $\alpha_\delta \leftarrow \sqrt{\frac{2\delta}{q_K}}$
 - 7: $\theta_{\text{new}} \leftarrow \theta_{\text{old}} + \alpha_\delta x_K$
 - 8: **return** θ_{new}
-

The conjugate-gradient solve produces

$$x_K \approx \left(\widehat{F}_{\text{old}} + \lambda_{\text{damp}} I \right)^{-1} \widehat{g}_{\text{old}},$$

while

$$q_K = x_K^\top \widehat{F}_{\text{old}} x_K$$

is computed using one additional undamped Fisher-vector product.

For fixed-learning-rate NPG, the trust-region scaling steps are omitted and the final update is simply

$$\theta_{\text{new}} = \theta_{\text{old}} + \alpha x_K.$$