

# MF-PhiBE for LQR Mean-Field Control from Discrete-Time Data

---

**Martin Hernandez**

UCLA, Department of Statistics & Data Science

Joint work with Yuhua Zhu, Qinxin Yan, and Erhan Bayraktar

`martinh@g.ucla.edu`

Foundations of Multi-Agent and Mean Field Reinforcement Learning

May 2026

The UCLA logo, consisting of the letters "UCLA" in white, bold, sans-serif font, centered within a solid blue rectangular background.

# Data and problem setting

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

Question. Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$V^*(\mu) = \mathbb{E}_{\mu_0 = \mu} \left[ \int_0^\infty e^{-\lambda t} r_\lambda^*(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^*(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$V^*(\mu) = \sup_{\pi} V^{\pi}(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top Q \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(s_t)$  and

$$ds_t = (As_t + \bar{A}\bar{\mu}_t + B a_t) dt + \sigma dW_t, \quad \bar{\mu}_t = \mathbb{E}[s_t].$$

Using Riccati.

$$\pi^*(\cdot \mid s, \mu) = \mathcal{N}\left(-N^{-1}B^\top(K(s - \bar{\mu}) + \Lambda\bar{\mu}), \frac{\lambda}{2}N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

The deterministic problem ( $\gamma = 0$ )  
defines the same optimal policy.

# Data and problem setting

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$V^*(\mu) = \mathbb{E}_{\pi^*} \left[ \int_0^\infty e^{-\lambda t} r^*(s, \mu, a) dt \mid \mu_0 = \mu \right],$$

where  $r^*(s, \mu) = \mathbb{E}_{\pi^*}[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$V^*(\mu) = \sup_{\pi} V^{\pi}(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^T Q s + \bar{\mu}^T Q \bar{\mu} + a^T N a)$$

where  $\mu_t = \text{Law}(x_t)$  and

$$dx_t = (Ax_t + A\bar{\mu}_t + B a_t) dt + \sigma dW_t, \quad \bar{\mu}_t = \mathbb{E}[x_t].$$

Using Riccati.

$$\pi^*(\cdot \mid s, \mu) = \mathcal{N}\left(-N^{-1}B^T(K(s-\bar{\mu}) + \Lambda\bar{\mu}), \frac{\lambda}{2}N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

The deterministic problem ( $\gamma = 0$ )  
defines the same optimal policy.

# Data and problem setting

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_a[r(s, \mu, a) - \lambda \log p_\pi(a)]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu).$$

where  $\pi(a_t | s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^T Q s + \beta^T Q \bar{\mu} + a^T N a)$$

where  $\mu_t = \text{Law}(x_t)$  and

$$dx_t = (Ax_t + A\bar{\mu}_t + B a_t) dt + \sigma dW_t, \quad \bar{\mu}_t = \mathbb{E}[x_t].$$

Using Riccati.

$$r^*(\cdot | s, \mu) = \mathcal{N}\left(-N^{-1}B^T(K(s - \bar{\mu}) + \Lambda\bar{\mu}), \frac{\lambda}{2}N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

The deterministic problem ( $\gamma = 0$ )  
defines the same optimal policy.

# Data and problem setting

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

---

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$

and  $\mu_t$  is given by

$$\dot{\mu}_t = A \mu_t + B \pi(s, \mu_t),$$

where  $\pi(s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^T Q s + \beta^T Q \bar{\mu} + a^T N a)$$

where  $\mu_t = \text{Law}(x_t)$  and

$$\dot{x}_t = (A x_t + A \bar{\mu}_t + B a_t) dt + \sigma dW_t, \quad \bar{\mu}_t = \mathbb{E}[x_t].$$

Using Riccati.

$$v^*(\cdot | s, \mu) = \mathcal{N}\left(-N^{-1} \Theta^T (K(s - \bar{\mu}) + \Lambda \bar{\mu}), \frac{1}{2} N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

The deterministic problem ( $\gamma = 0$ ) defines the same optimal policy.

# Data and problem setting

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(s, \mu)$  is a random and feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^T Q s + \beta^T Q \bar{\mu} + a^T N a)$$

where  $\mu_t = \text{Law}(x_t)$  and

$$dx_t = (Ax_t + A\bar{\mu}_t + B a_t) dt + \sigma dW_t, \quad \bar{\mu}_t = \mathbb{E}[x_t].$$

Using Riccati.

$$v^*(\cdot | s, \mu) = \mathcal{N}\left(-N^{-1}B^T(K(s-\bar{\mu}) + \Lambda\bar{\mu}), \frac{1}{2}N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

The deterministic problem ( $\gamma = 0$ )  
defines the same optimal policy.

# Data and problem setting

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

---

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

---

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top \bar{Q} \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(x_t)$  and  $\bar{\mu} = \begin{bmatrix} \mu \\ \dot{\mu} \end{bmatrix}$ ,  $\dot{\mu} = (A\mu + B\dot{\mu} + D\dot{a})dt + \sigma dW_t$ ,  $A_t = A(s)$ .

Using Riccati:  
$$v^*(\cdot | \cdot, \mu) = \mathcal{N}\left(-N^{-1} \Theta \left(K(-\beta) + \Lambda D\right), \frac{\beta}{2} N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

The deterministic problem ( $\gamma = 0$ ) defines the same optimal policy.



Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top \bar{Q} \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(s_t)$  and

$$ds_t = (As_t + \bar{A}\bar{\mu}_t + Ba_t) dt + \gamma dW_t, \quad \bar{\mu}_t = \mathbb{E}[s_t].$$

Using Riccati:

$$v^*(s, \mu) = N \left( -N^{-1} \bar{Q} (K(-\beta) + \Lambda) + \frac{\beta}{2} N^{-1} \right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

Remark:

The deterministic problem ( $\gamma = 0$ ) defines the same optimal policy.

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top \bar{Q} \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(s_t)$  and

$$ds_t = (As_t + \bar{A}\bar{\mu}_t + Ba_t) dt + \gamma dW_t, \quad \bar{\mu}_t = \mathbb{E}[s_t].$$

Using Riccati.

$$V^*(s, \mu) = \mathcal{N}\left(-N^{-1} \bar{Q} (K(-\beta) + \Lambda D) \frac{\gamma}{2} N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

Use deterministic problem ( $\gamma = 0$ )  
defines the same optimal policy

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top \bar{Q} \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(s_t)$  and

$$ds_t = (As_t + \bar{A}\bar{\mu}_t + Ba_t) dt + \gamma dW_t, \quad \bar{\mu}_t = \mathbb{E}[s_t].$$

Using Riccati.

$$\pi^*(\cdot \mid s, \mu) = \mathcal{N}\left(-N^{-1}B^\top(K(s - \bar{\mu}) + \Lambda\bar{\mu}), \frac{\lambda}{2}N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation  
independent of  $\gamma$ .

Then the LQR problem (1) and (2)  
define the same optimal policy.

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$   
and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top \bar{Q} \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(s_t)$  and

$$ds_t = (As_t + \bar{A}\bar{\mu}_t + Ba_t) dt + \gamma dW_t, \quad \bar{\mu}_t = \mathbb{E}[s_t].$$

Using Riccati.

$$\pi^*(\cdot \mid s, \mu) = \mathcal{N}\left(-N^{-1}B^\top(K(s - \bar{\mu}) + \Lambda\bar{\mu}), \frac{\lambda}{2}N^{-1}\right)$$

where  $K$  and  $\Lambda$  solve a Riccati equation independent of  $\gamma$ .

Can the continuous problem (and  $\gamma$ )  
define the same optimal policy?

Observed data at step  $\Delta t$ .

$$\mathcal{D} = \{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}.$$

**Question.** Can one use one-step data while keeping the continuous-time control problem, rather than replacing it by a discrete-time Bellman equation?

Consider the value function

$$\mathcal{V}^\pi(\mu) = \mathbb{E}_{s \sim \mu} \left[ \int_0^\infty e^{-\beta t} r_\lambda^\pi(s, \mu_t) dt \mid \mu_0 = \mu \right],$$

where  $r_\lambda^\pi(s, \mu) = \mathbb{E}_\pi[r(s, \mu, a) - \lambda \log p_\pi]$  and denote by

$$\mathcal{V}^*(\mu) = \sup_{\pi} \mathcal{V}^\pi(\mu),$$

where  $\pi(da \mid s, \mu)$  is a randomized feedback policy.

LQR benchmark.

$$r(s, \mu, a) = -(s^\top Q s + \bar{\mu}^\top \bar{Q} \bar{\mu} + a^\top N a)$$

where  $\mu_t = \text{Law}(s_t)$  and

$$ds_t = (A s_t + \bar{A} \bar{\mu}_t + B a_t) dt + \gamma dW_t, \quad \bar{\mu}_t = \mathbb{E}[s_t].$$

Using Riccati.

$$\pi^*(\cdot \mid s, \mu) = \mathcal{N}\left(-N^{-1} B^\top (\textcolor{blue}{K}(s - \bar{\mu}) + \textcolor{blue}{\Lambda} \bar{\mu}), \frac{\lambda}{2} N^{-1}\right)$$

where  $\textcolor{blue}{K}$  and  $\textcolor{blue}{\Lambda}$  solve a Riccati equation independent of  $\gamma$ .

Key Message

The deterministic problem ( $\gamma = 0$ ) defines the same optimal policy.

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(C_{1,x}^\pi - \beta)V^\pi(\mu) + r^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^2 - s}{\Delta t} \mid s, \mu, a\right],$$

**LQR structure.** For the LQR benchmark,

$$b(s, \mu, a) = As + \bar{A}\bar{\mu} + \bar{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(C_{1,p}^\pi - \beta)\hat{V}^\pi(\mu) + r^\pi(\mu) = 0,$$

and

$$\hat{V}^*(\mu) = \sup_a \hat{V}^a(\mu).$$

## Optimal MF-PhiBE policy

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\lambda}{2} N^{-1})$ , with

$$\hat{m}^*(s, \mu) = -N^{-1}\hat{B}^\top(K(s - \bar{\mu}) + \bar{A}\bar{\mu}), \quad m^*(s, \mu) = -N^{-1}B^\top(K(s - \bar{\mu}) + \bar{A}\bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{s,\mu}\Delta t, \quad |V^{\pi^*} - \hat{V}^{\hat{\pi}^*}| \leq C_\mu(\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\hat{\pi}^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) = \mathbb{E}\left[\frac{s_t^\Sigma - s}{\Delta t} \mid s, \mu, a\right],$$

LQR structure. For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + \hat{r}_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_a \hat{\mathcal{V}}^a(\mu).$$

## MF-PhiBE with LQR structure

The optimal MF-PhiBE policy is Gaussian,  $\hat{\nu}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\lambda}{2} N^{-1})$ , with

$$\hat{m}^*(s, \mu) = -N^{-1} \hat{B}^\top (\hat{K}(s - \bar{\mu}) + \hat{\Lambda} \bar{\mu}), \quad m^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{b,\mu} \Delta t, \quad |\mathcal{V}^{\hat{\nu}^*} - \mathcal{V}^{\nu^*}| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\hat{\nu}^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) = \mathbb{E}\left[\frac{s_t^2 - s}{\Delta t} \mid s, \mu, a\right],$$

LQR structure. For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + \hat{r}_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_a \hat{\mathcal{V}}^a(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\mathcal{V}}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\lambda}{2} N^{-1})$ , with

$$\hat{m}^*(s, \mu) = -N^{-1} \hat{B}^T (K(s - \bar{\mu}) + \Lambda \bar{\mu}), \quad m^*(s, \mu) = -N^{-1} B^T (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{b,\mu} \Delta t, \quad |\hat{\mathcal{V}}^* - \mathcal{V}^*| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\hat{\mathcal{V}}^*(s, \mu)]$ .



# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^+ - s}{\Delta t} \mid s, \mu, a\right].$$

**LQR structure.** For the LQR benchmark,

$$b(s, \mu, a) = \bar{A}s + \bar{A}\bar{\mu} + \bar{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + \hat{r}_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_a \hat{\mathcal{V}}^a(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\lambda}{2} N^{-2})$ , with

$$\hat{m}^*(s, \mu) = -N^{-1} \bar{B}^\top (\bar{K}(s - \bar{\mu}) + \bar{\Lambda} \bar{\mu}), \quad m^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{b,\mu} \Delta t, \quad |\mathcal{V}^{\hat{\pi}^*} - \mathcal{V}^{*\pi}| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[m^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^+ - s}{\Delta t} \mid s, \mu, a\right].$$

**LQR structure.** For the LQR benchmark,

$$b(s, \mu, a) = \bar{A}s + \bar{A}_\mu\mu + \bar{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + r_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_a \hat{\mathcal{V}}^a(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\lambda}{2} N^{-1})$ , with

$$\hat{m}^*(s, \mu) = -N^{-1} \bar{B}^\top (K(s - \bar{\mu}) + \bar{\Lambda} \bar{\mu}), \quad m^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{b,\mu} \Delta t, \quad |\hat{\mathcal{V}}^{\pi^*} - \mathcal{V}^{\pi^*}| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[m^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^+ - s}{\Delta t} \mid s, \mu, a\right].$$

**LQR structure.** For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_a \hat{\mathcal{V}}^*(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{\pi}^*(s, \mu), \frac{\Delta t}{2} N^{-1})$ , with

$$\hat{\pi}^*(s, \mu) = -N^{-1} \hat{B}^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}), \quad \pi^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{\pi}^*(s, \mu) - \pi^*(s, \mu)\| \leq C_{\hat{\pi},\mu} \Delta t, \quad |\mathcal{V}^{\hat{\pi}^*} - \mathcal{V}^{\pi^*}| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\pi^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\hat{\pi}^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^+ - s}{\Delta t} \mid s, \mu, a\right].$$

**LQR structure.** For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_a \hat{\mathcal{V}}^*(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\Delta t}{2} N^{-2})$ , with

$$\hat{m}^*(s, \mu) = -N^{-1} \hat{B}^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}), \quad m^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{b,\mu} \Delta t, \quad \|\mathcal{V}^{\hat{\pi}^*} - \mathcal{V}^{*\pi}\| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[m^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^+ - s}{\Delta t} \mid s, \mu, a\right].$$

**LQR structure.** For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{\hat{b},0}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + r_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^\pi(\mu) = \sup_a \hat{\mathcal{V}}^a(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{\pi}^*(s, \mu), \frac{\beta}{2} N^{-2})$ , with

$$\hat{\pi}^*(s, \mu) = -N^{-1} \hat{B}^\top (K(s - \bar{p}) + \Lambda \bar{p}), \quad \pi^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{p}) + \Lambda \bar{p}),$$

then

$$\|\hat{\pi}^*(s, \mu) - \pi^*(s, \mu)\| \leq C_{\hat{b},\mu} \Delta t, \quad |\mathcal{V}^{\pi^*} - \hat{\mathcal{V}}^{\hat{\pi}^*}| \leq C_{\hat{b}} (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\pi^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\hat{\pi}^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E}\left[\frac{s^+ - s}{\Delta t} \mid s, \mu, a\right].$$

**LQR structure.** For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{\hat{b},0}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + r_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_{\pi} \hat{\mathcal{V}}^\pi(\mu).$$

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot \mid s, \mu) = \mathcal{N}(\hat{\pi}^*(s, \mu), \frac{1}{2}N^{-1})$ , with

$$\hat{\pi}^*(s, \mu) = -N^{-1}\hat{B}^\top(K(s - \bar{p}) + \Lambda\bar{p}), \quad \pi^*(s, \mu) = -N^{-1}B^\top(K(s - \bar{p}) + \Lambda\bar{p}),$$

then

$$\|\hat{\pi}^*(s, \mu) - \pi^*(s, \mu)\| \leq C_{\hat{b},\mu}\Delta t, \quad |\mathcal{V}^{\hat{\pi}^*} - \mathcal{V}^{\pi^*}| \leq C_{\hat{b}}(\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[\hat{\pi}^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\pi^*(s, \mu)]$ .

# MF-PhiBE and LQR policy approximation

**Policy evaluation.** For a fixed policy, the exact continuous-time value solves

$$(\mathcal{L}_{b,\Sigma}^\pi - \beta)\mathcal{V}^\pi(\mu) + r_\lambda^\pi(\mu) = 0.$$

MF-PhiBE keeps the continuous-time equation and replaces the unknown local drift by

$$\hat{b}(s, \mu, a) := \mathbb{E} \left[ \frac{s^+ - s}{\Delta t} \middle| s, \mu, a \right].$$

**LQR structure.** For the LQR benchmark,

$$\hat{b}(s, \mu, a) = \hat{A}s + \hat{A}\bar{\mu} + \hat{B}a.$$

Since  $\gamma = 0$  defines the same policy, MF-PhiBE

$$(\mathcal{L}_{\hat{b},0}^\pi - \beta)\hat{\mathcal{V}}^\pi(\mu) + r_\lambda^\pi(\mu) = 0,$$

and

$$\hat{\mathcal{V}}^*(\mu) = \sup_{\pi} \hat{\mathcal{V}}^\pi(\mu).$$

## MF-PhiBE error in LQR

The optimal MF-PhiBE policy is Gaussian,  $\hat{\pi}^*(\cdot | s, \mu) = \mathcal{N}(\hat{m}^*(s, \mu), \frac{\lambda}{2} N^{-1})$ , with

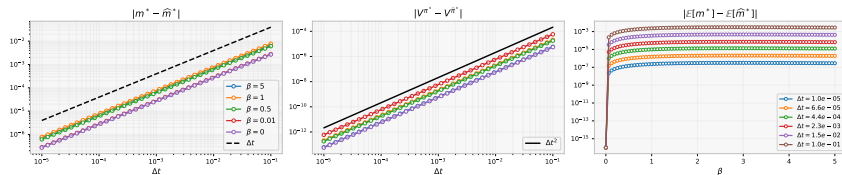
$$\hat{m}^*(s, \mu) = -N^{-1} \hat{B}^\top (\hat{K}(s - \bar{\mu}) + \hat{\Lambda} \bar{\mu}), \quad m^*(s, \mu) = -N^{-1} B^\top (K(s - \bar{\mu}) + \Lambda \bar{\mu}),$$

then

$$\|\hat{m}^*(s, \mu) - m^*(s, \mu)\| \leq C_{s,\mu} \Delta t, \quad |\mathcal{V}^{\pi^*} - \mathcal{V}^{\hat{\pi}^*}| \leq C_\mu (\Delta t)^2.$$

Moreover, when  $\beta = 0$ , we have  $\mathbb{E}_{s \sim \mu}[m^*(s, \mu)] = \mathbb{E}_{s \sim \mu}[\hat{m}^*(s, \mu)]$ .

## Convergence order verification:

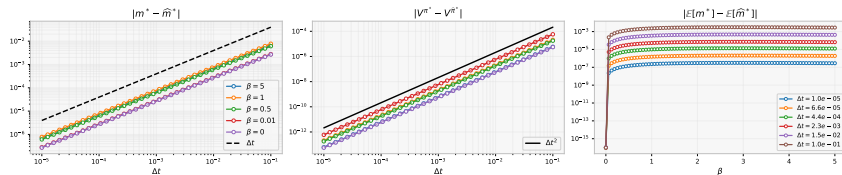


## Algorithm for LQR: Levenberg-Marquardt





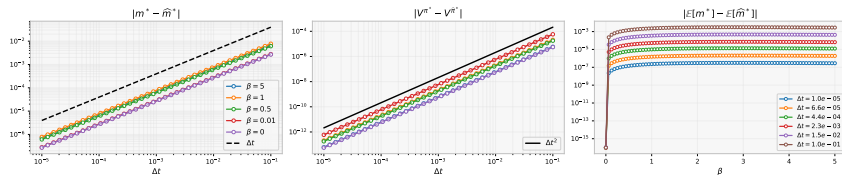
## Convergence order verification:



## Algorithm for LQR:



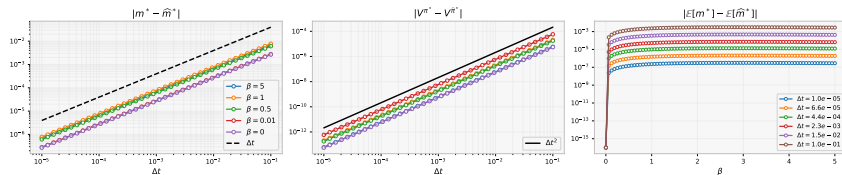
## Convergence order verification:



Algorithm for LQR: Let  $\pi_\omega$ .



## Convergence order verification:



Algorithm for LQR: Let  $\pi_\omega$ .

DATA

$$\{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}$$

Estimate  $\Delta t$ :

$$\Delta t = \frac{1}{N_D} \sum_{\eta=1}^{N_D} \frac{1}{\mu_\eta}$$

Estimate  $\Delta t$ :

$$\Delta t = \frac{1}{N_D} \sum_{\eta=1}^{N_D} \frac{1}{\mu_\eta}$$

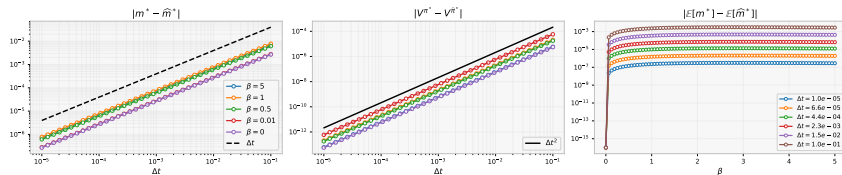
Policy gradient:

$$\nabla_{\omega} J(\omega) = \frac{1}{N_D} \sum_{\eta=1}^{N_D} \frac{1}{\mu_\eta} \nabla_{\omega} \log \mu_\eta (s_\eta, a_\eta)$$

Estimate  $\Delta t$ :

$$\Delta t = \frac{1}{N_D} \sum_{\eta=1}^{N_D} \frac{1}{\mu_\eta}$$

## Convergence order verification:



Algorithm for LQR: Let  $\pi_\omega$ .

DATA

$$\{(s_\eta, \mu_\eta, a_\eta, s_\eta^+, r_\eta)\}_{\eta=1}^{N_D}$$

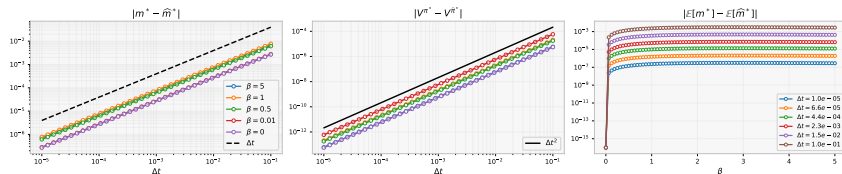
Estimate drift:

$$\hat{b} \sim \tilde{b}_\eta := \frac{s_\eta^+ - s_\eta}{\Delta t}$$

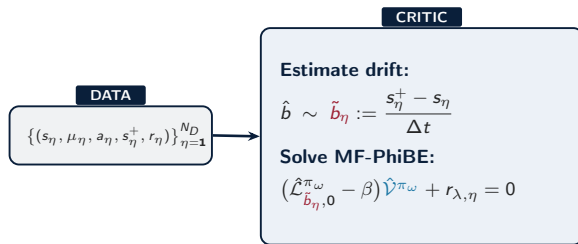
Policy gradient:

$$\nabla_{\theta} J(\theta) = \mathbb{E} \left[ \sum_{\eta=1}^{N_D} \tilde{b}_\eta \nabla_{\theta} \pi_\omega(s_\eta, a_\eta) \right]$$

## Convergence order verification:

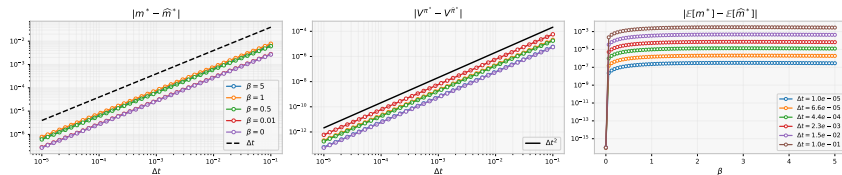


Algorithm for LQR: Let  $\pi_\omega$ .

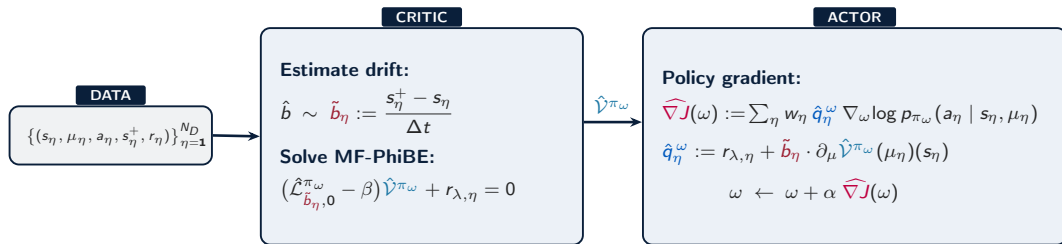


# Model-free actor-critic

## Convergence order verification:

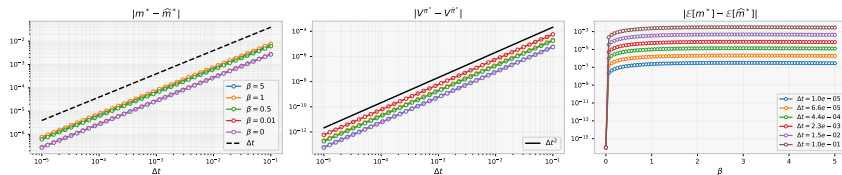


## Algorithm for LQR: Let $\pi_\omega$ .

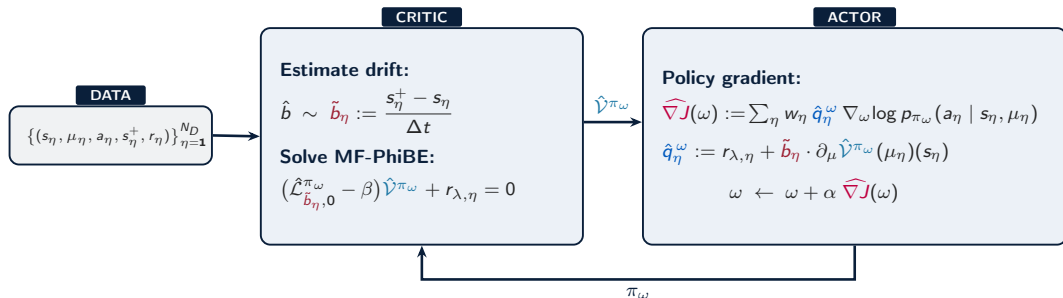


# Model-free actor-critic

## Convergence order verification:



## Algorithm for LQR: Let $\pi_\omega$ .



- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s + \mu) + \omega_2 \bar{\mu}, \frac{1}{2} N^{-1})$$

$$V_{\theta}(\mu) = -(\theta_0 + \theta_1 \bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhiBE with Galerkin with cylindrical basis functions.
- Only the first and second moments of  $\mu$  are needed.

Beyond vanilla gradient descent: Since we have access to  $\widehat{\nabla J}$ , we can implement different gradient-based optimization methods.



- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s + \mu) + \omega_2\bar{\mu}, \frac{1}{2}N^{-1})$$

$$V_{\theta}(\mu) = -(\theta_0 + \theta_1\bar{\mu} + \theta_2\text{Var}(\mu))$$

- We solve PhIBE with Galerkin with cylindrical basis functions.
- Only the first and second moments of  $\mu$  are needed.

Beyond vanilla gradient descent: Since we have access to  $\nabla J$ , we can implement different gradient-based optimization methods.

- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2 \bar{\mu}, \frac{\lambda}{2} N^{-1})$$

$$V_{\theta}(\mu) = -(\theta_0 + \theta_1 \bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhBE with Galerkin with cylindrical basis functions.
- Only the first and second moments of  $\mu$  are needed.

Beyond vanilla gradient descent: Since we have access to  $\nabla J$ , we can implement different gradient-based optimization methods.

- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2\bar{\mu}, \frac{\lambda}{2}N^{-1})$$

$$\mathcal{V}_{\theta}(\mu) = -(\theta_0 + \theta_1\bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhIDE with Galerkin with cylindrical basis functions.
- Only the first and second moments of  $\mu$  are needed.

Beyond vanilla gradient descent: Since we have access to  $\nabla J$ , we can implement different gradient-based optimization methods.

- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2 \bar{\mu}, \frac{\lambda}{2} N^{-1})$$

$$\mathcal{V}_{\theta}(\mu) = -(\theta_0 + \theta_1 \bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhiBE with Galerkin with cylindrical basis functions.

► Only the first and second moments of  $\mu$  are needed.

Beyond vanilla gradient descent: Since we have access to  $\nabla J$ , we can implement different gradient-based optimization methods.

- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2 \bar{\mu}, \frac{\lambda}{2} N^{-1})$$

$$\mathcal{V}_{\theta}(\mu) = -(\theta_0 + \theta_1 \bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhiBE with Galerkin with cylindrical basis functions.
- **Only the first and second moments of  $\mu$  are needed.**

Beyond vanilla gradient descent: Since we have access to  $\nabla J$ , we can implement different gradient-based optimization methods.

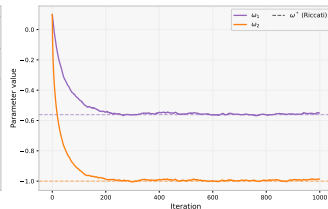
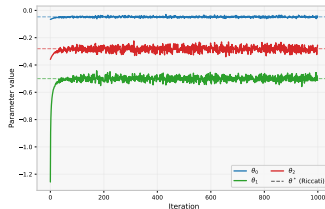
# Numerical implementation

- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2\bar{\mu}, \frac{\lambda}{2}N^{-1})$$

$$\mathcal{V}_{\theta}(\mu) = -(\theta_0 + \theta_1\bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhiBE with Galerkin with cylindrical basis functions.
- **Only the first and second moments of  $\mu$  are needed.**



Beyond vanilla gradient descent: Since we have access to  $\nabla J$ , we can implement different gradient-based optimization methods.

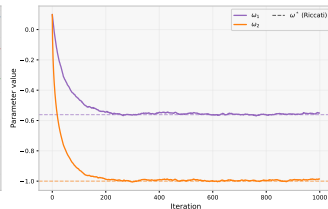
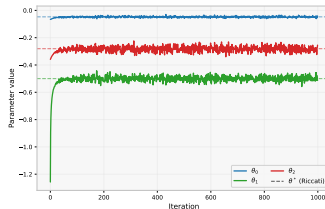
# Numerical implementation

- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2\bar{\mu}, \frac{\lambda}{2}N^{-1})$$

$$\mathcal{V}_{\theta}(\mu) = -(\theta_0 + \theta_1\bar{\mu} + \theta_2 \text{Var}(\mu))$$

- We solve PhiBE with Galerkin with cylindrical basis functions.
- **Only the first and second moments of  $\mu$  are needed.**



**Beyond vanilla gradient descent:** Since we have access to  $\widehat{\nabla} J$ , we can implement different gradient-based optimization methods.

# Numerical implementation

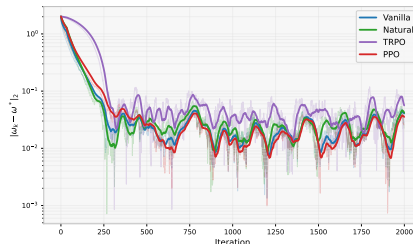
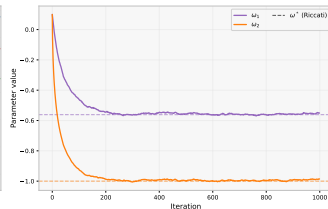
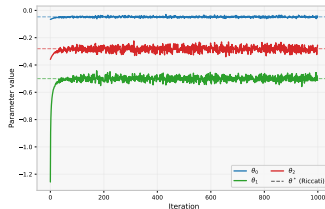
- The value function and the policy are parameterized:

$$\pi_{\omega}(\cdot|s, \mu) = \mathcal{N}(\omega_1(s - \bar{\mu}) + \omega_2 \bar{\mu}, \frac{\lambda}{2} N^{-1})$$

$$\mathcal{V}_{\theta}(\mu) = -(\theta_0 + \theta_1 \bar{\mu} + \theta_2 \text{Var}(\mu))$$

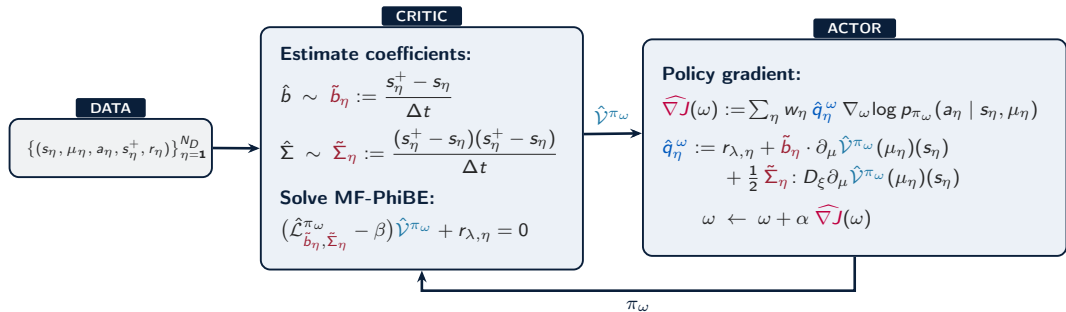
- We solve PhiBE with Galerkin with cylindrical basis functions.
- Only the first and second moments of  $\mu$  are needed.

**Beyond vanilla gradient descent:** Since we have access to  $\widehat{\nabla J}$ , we can implement different gradient-based optimization methods.



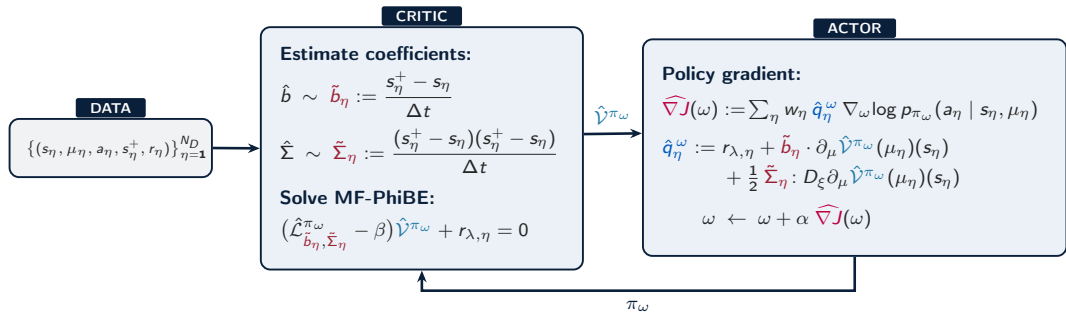


# Model-free actor-critic: general setting



Thanks for your attention!

# Model-free actor-critic: general setting



Thanks for your attention!

# Modeling Stochastic Multi-Agent Interaction in Intraday Battery Energy Storage Dispatch with Market Power

Hezhong Zhang

Department of Statistics and Applied Probability  
University of California, Santa Barbara

May 2026

<https://arxiv.org/pdf/2605.01178>

Joint work with Ruimeng Hu and Mike Ludkovski

# Stochastic Differential Game Model ( $N$ agents)

## Exogenous supply dynamics

$$dQ_t = \kappa(\theta(t) - Q_t) dt + \sigma^0 dW_t^0 \quad (1)$$

## Battery State of Charge dynamics

$$dS_t^i = (\mu_t^i + \alpha_t^i) dt + \sigma^i \rho^i dW_t^0 + \sigma^i \sqrt{1 - (\rho^i)^2} d\tilde{W}_t^i, \quad (2)$$

where  $\mu_t^i = a^i(t) + b^i(t)Q_t$  is hybrid asset self generation, and  $\alpha_t^i = \alpha^i(t, Q_t, \mathbf{S}_t)$ , charges ( $\alpha_t^i > 0$ ) and discharges ( $\alpha_t^i < 0$ ).

## Price

$$P_t^i = \underbrace{\bar{P}^i}_{\text{Base price}} - c_1^i \underbrace{\left( Q_t - \mathbf{w}_i^\top \alpha_t \right)}_{\text{Net Supply}} \quad (3)$$

## Optimization problem (Linear Quadratic (LQ) Structure)

$$\hat{\alpha}^i \in \arg \min_{\alpha^i} J^i(\alpha^i, \hat{\alpha}^{-i}),$$

$$J^i(\alpha) = \mathbb{E} \left[ \int_0^T \underbrace{P_t^i \alpha_t^i + c_2^i (\alpha_t^i)^2 + c_3^i (S_t^i - \zeta^i(t))^2}_{f^i(t, Q_t, \mathbf{S}_t, \alpha_t^i, \alpha_t^{-i})} dt + \underbrace{c_4^i (S_T^i - \zeta^i(T))^2}_{g^i(Q_t, \mathbf{S}_t)} \right] \quad (4)$$

# Main Technical Contribution

Markovian LQ Differential Game  $\implies$  Best Response Functions  $\implies$   
Coupled HJB System  $\implies$  LQ Ansatz  $\implies$  Coupled Riccati ODEs

- The best response function  $V^i$  for each agent  $i$  evaluated at Nash Equilibrium  $\hat{\alpha}$  satisfies

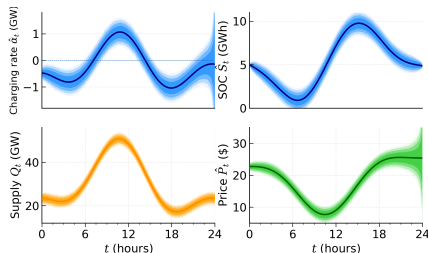
$$\partial_t V^i(t, q, \mathbf{s}) + \inf_{\alpha^i} \left\{ f^i(t, q, \mathbf{s}, \alpha^i, \hat{\alpha}^{-i}) + \mathcal{L}^{\alpha^i, \hat{\alpha}^{-i}} V^i(t, q, \mathbf{s}) \right\} = 0, \quad (5)$$

with  $V^i(T, q, \mathbf{s}) = g^i(q, \mathbf{s})$ . The LQ ansatz is of the form.

$$V^i(t, q, \mathbf{s}) = \begin{bmatrix} q \\ \mathbf{s} \end{bmatrix}^\top \begin{bmatrix} p_0^i(t) & \mathbf{p}^i(t)^\top \\ \mathbf{p}^i(t) & \mathbf{P}^i(t) \end{bmatrix} \begin{bmatrix} q \\ \mathbf{s} \end{bmatrix} + \begin{bmatrix} r_0^i(t) \\ \mathbf{r}^i(t) \end{bmatrix}^\top \begin{bmatrix} q \\ \mathbf{s} \end{bmatrix} + u^i(t). \quad (6)$$

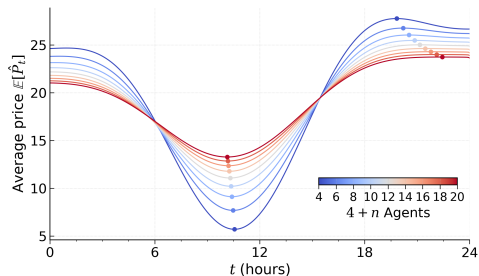
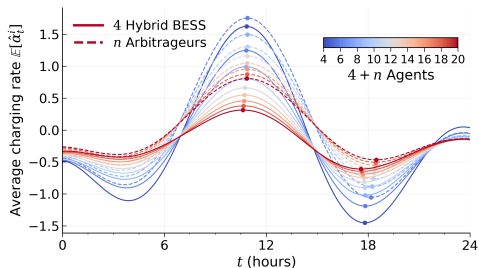
- The heterogeneous setting requires solving  $\mathcal{O}(N^3)$  coupled Riccati ODEs, while under homogeneous setting the system collapses to  $\mathcal{O}(1)$ .

# What Does Equilibrium Look Like? (Homogeneous, $N = 8$ )



- Batteries charge when prices are low, and discharge during high-price periods
- SOC decreases in the discharging period and increase during charging
- Resulting prices are lowest in the morning and highest during the evening.
- Evening Prices flattened

# How Does Competition Reshape Equilibrium?



- More competition reduces individual market power, leading to less charging and discharging
- Gap between the charging rate of Hybrid and Arbitrageurs agents increase as number of agents increase
- Peak shifts down and valley shifts up, meaning price goes up in the charging period and goes down when discharging

## Main Takeaways

- Strategic BESS interaction creates endogenous price effects
- Increasing competition reduces individual market power
- The LQ game framework enables tractable equilibrium analysis

## Future Directions

- Non-LQ structure
- Deep learning / reinforcement learning methods for non-LQ games
- Major–Minor and mean field game formulations
- Heterogeneous and asymmetric equilibrium structures

**Thank you!**



# Population-Aware Imitation Learning in Mean-Field Games with Common Noise

Grégoire Lambrecht

Joint work with Mathieu Laurière

NYU Shanghai

IMSI Workshop, Chicago

May 2026



### Key Contributions of Lambrecht and Laurière (2026)

- ▶ **Theory:** Convergence guarantees for imitation learning proxies in Mean-Field Games with common noise.
- ▶ **Algorithm:** *Deep Fictitious Play*, a heuristic extension of *Fictitious Play* to the common noise setting.
- ▶ **Robustness:** Empirical evidence showing the advantage of *population-aware* policies over *population-unaware* ones.

arXiv preprint →



## What is Imitation Learning in MFGs?

- ▶ Observe experts playing at equilibrium.
- ▶ Assume that the agents' objective (i.e., the reward function) is unknown.
- ▶ Questions:
  - (Q1) **Equilibrium Recovery:** *Is it possible to recover actions to reconstruct an equilibrium ?*
  - (Q2) **Performance Matching:** *Is it possible to recover actions to perform well against the experts?*
- ▶ Without common noise: [Ramponi et al. \(2023\)](#).

Giorgia Ramponi, Pavel Kolev, Olivier Pietquin, Niao He, Mathieu Laurière, and Matthieu Geist. On imitation in mean-field games. *Advances in Neural Information Processing Systems*, 36, 2023.

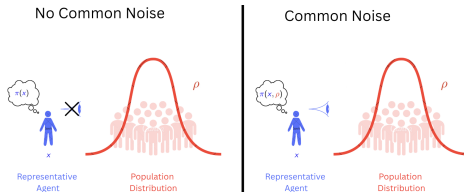
## What is common noise?

- An additional source of randomness affecting all agents simultaneously.



## What changes with common noise?

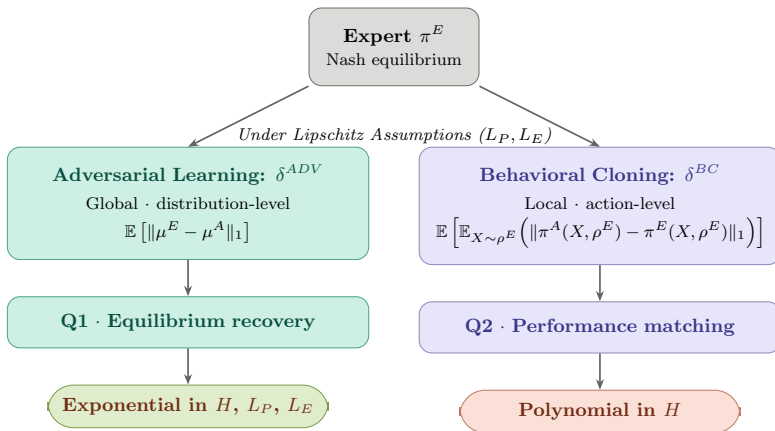
- The population distribution becomes stochastic.
- Policies become population-aware.



**Can population-unaware imitation learning methods still perform well under common noise?**

## Theoretical Results

- $\pi^E$ : Expert Policy,  $\pi^A$ : Candidate Policy.
- **(Q1) Equilibrium Recovery:**  $\text{Exploitability}(\pi^A) \leq C\delta(\pi^A, \pi^E)$
- **(Q2) Perf. Matching:**  $\text{Perf. vs. Expert}(\pi^A) \leq C'\delta(\pi^A, \pi^E)$



# Learning Algorithm, Numerical Results and Key Takeaways

Expert World  
at Equilibrium

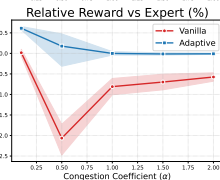
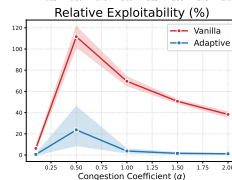
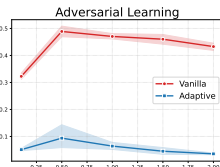
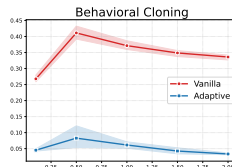


Observe  $(x_i, a_i)$

Learner World

Learning  $\pi^\theta(a|x, \rho)$

$$L(\theta) \sim \sum_{i=1}^4 \|\pi^\theta(x_i, \hat{p}) - \hat{\pi}(x_i)\|_1$$



Thank You

# Policy Gradient Learning for Distributionally Robust Markov Decision Processes under Wasserstein Ambiguity

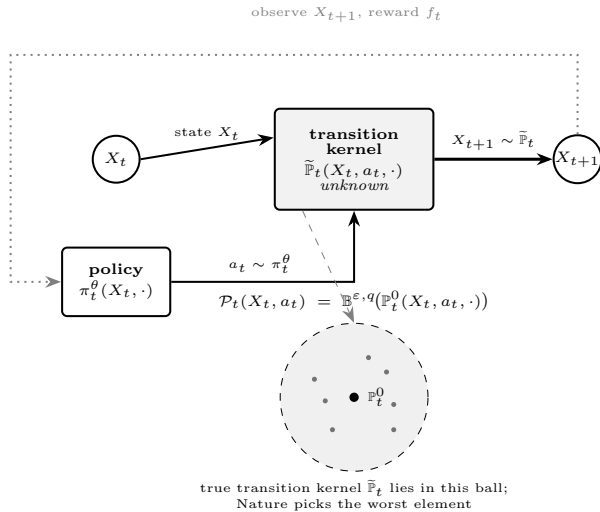
Y. Hafsi   S. Mekkaoui   H. Pham   K. Yan

CMAP, École Polytechnique  
School of Mathematical Sciences, Xiamen University

*Foundations of Multi-Agent and Mean Field Reinforcement Learning*

May 21, 2026

# From MDPs to Robust MDPs



## Standard MDP. Tuple

$(\mathcal{X}, A, \{\mathbb{P}_t\}_{t < T}, f, g, T)$  with  $\mathcal{X} \subset \mathbb{R}^d$ ,  $A \subset \mathbb{R}^m$  compact,  $T \in \mathbb{N}^*$ ,  $\mathcal{T} := \{0, \dots, T-1\}$ ,  $\mathbb{P}_t : \mathcal{X} \times A \rightarrow \mathcal{P}(\mathcal{X})$ ,  $f : \mathcal{T} \times \mathcal{X} \times A \times \mathcal{X} \rightarrow \mathbb{R}$ ,  $g : \mathcal{X} \rightarrow \mathbb{R}$ , and trajectory reward  $R(S) = \sum_t f(t, X_t, a_t, X_{t+1}) + g(X_T)$ .

## Randomized parametric policy.

$\pi_t^\theta : \mathcal{X} \rightarrow \mathcal{P}(A)$ ,  $\theta \in \Theta \subset \mathbb{R}^{d_\theta}$  compact; for  $\mu_0 \in \mathcal{P}(\mathcal{X})$ , optimize  $J(\theta) := \mathbb{E}_{\mu_0}[V_0^{\pi^\theta}(X_0)]$ .

## Specialization.

For  $\varepsilon \geq 0$ ,  $q \geq 1$ ,  $\mathbb{P}_t^0 : \mathcal{X} \times A \rightarrow \mathcal{P}(\mathcal{X})$ :  $\mathcal{P}_t(x, a) = \mathbb{B}^{\varepsilon, q}(\mathbb{P}_t^0(x, a, \cdot))$ ,  $\mathcal{W}_q$ -ball of radius  $\varepsilon$ .

## Robust criterion:

$$V(x) = \sup_{\pi} \inf_{\mathbb{P} \in \mathcal{B}_x, \pi} \mathbb{E}^{\mathbb{P}}[R(S)].$$



# Dynamic programming and the differentiation obstacle

## Theorem (Robust DPP, Thm. 2.1)

For all  $(t, x) \in \mathcal{T} \times \mathcal{X}$  and all admissible  $\pi$ ,

$$V_T^\pi(x) = g(x), \quad V_t^\pi(x) = \mathbb{E}_{a \sim \pi_t(x, \cdot)} \left[ \inf_{\mathbb{P}_t \in \mathcal{P}_t(x, a)} \mathbb{E}^{\mathbb{P}_t} [f(t, x, a, X_{t+1}) + V_{t+1}^\pi(X_{t+1})] \right].$$

A Borel-measurable worst-case selector  $\mathbb{P}_t^{*, \pi}$  attains the infimum and  $V^\pi(x_0) = \mathbb{E}^{\mathbb{P}^{*, \pi}}[R(S)]$ .

**Obstacle.** The recursion under the parametric policy reads

$$V_t^\theta(x) = \int_A \int_{\mathcal{X}} [f(t, x, a, x') + V_{t+1}^\theta(x')] \mathbb{P}_t^{*, \theta}(x, a, dx') \pi_t^\theta(x, da).$$

The kernel  $\mathbb{P}_t^{*, \theta}$  depends implicitly on  $\theta$  through  $V_{t+1}^\theta$ . In general  $\theta \mapsto \mathbb{P}_t^{*, \theta}$  admits no closed form and is not amenable to a straightforward induction.

**Strategy.** Pass to a dual representation in which  $\mathbb{P}_t^{*, \theta}$  never appears explicitly.

# Wasserstein duality and policy gradient

**Constrained OT.** For  $V : \mathcal{X} \rightarrow \mathbb{R}$  continuous,  $\mathbb{P}^0 \in \mathcal{P}(\mathcal{X})$ ,  $\varepsilon \geq 0$ ,  $c(x, y) = \|x - y\|^q$ :

$$\inf_{\mathbb{P} \in \mathbb{B}^{\varepsilon, q}(\mathbb{P}^0)} \mathbb{E}^{\mathbb{P}}[V(Y)] = \inf_{\substack{\gamma \in \mathcal{P}(\mathcal{X} \times \mathcal{X}): \\ \text{pr}_{1\#} \gamma = \mathbb{P}^0, \mathbb{E}^{\gamma}[c(X, Y)] \leq \varepsilon^q}} \mathbb{E}^{\gamma}[V(Y)].$$

**Strong duality** (Blanchet–Murthy, 2019). With  $\mathcal{F}^\lambda(V)(x) := \sup_{y \in \mathcal{X}} \{V(y) - \lambda c(x, y)\}$ ,  $\lambda \geq 0$ :

$$\inf_{\mathbb{P} \in \mathbb{B}^{\varepsilon, q}(\mathbb{P}^0)} \mathbb{E}^{\mathbb{P}}[V(Y)] = \sup_{\lambda \geq 0} \mathbb{E}^{\mathbb{P}^0} \left[ \inf_{y \in \mathcal{X}} \{V(y) + \lambda c(X, y)\} \right] - \varepsilon^q \lambda.$$

**Dual robust Bellman.** For  $(t, \theta, x, a) \in \mathcal{T} \times \Theta \times \mathcal{X} \times A$ :

$$G_t^\theta(x, a) := \sup_{\lambda \geq 0} \left\{ \mathbb{E}_{X \sim \mathbb{P}_t^0(x, a, \cdot)} \left[ -\mathcal{F}^\lambda(-f(t, x, a, \cdot) - V_{t+1}^\theta)(X) \right] - \varepsilon^q \lambda \right\} \in \mathbb{R},$$

$$\Lambda_t^{*, \theta}(x, a) := \arg \max_{\lambda \in [0, \Lambda]} \{ \dots \} \subset [0, \Lambda], \quad Y_{t, \theta, \lambda}^*(x, a; X) := \operatorname{argmin}_{y \in \mathcal{X}} \{ f + V_{t+1}^\theta + \lambda c(X, y) \} \subset \mathcal{X}.$$

## Theorem (Policy gradient, Thm. 3.1)

For every  $r \in \mathbb{R}^{d_\theta}$ ,  $\theta \mapsto G_t^\theta(x, a)$  admits one-sided directional derivatives:

$$\begin{aligned} D_\theta^+ G_t^\theta(x, a)[r] &= \sup_{\lambda \in \Lambda_t^{*, \theta}(x, a)} \mathbb{E}_{X \sim \mathbb{P}_t^0(x, a, \cdot)} \left[ \inf_{y \in Y_{t, \theta, \lambda}^*(x, a; X)} D_\theta V_{t+1}^\theta(y)[r] \right], \\ D_\theta^- G_t^\theta(x, a)[r] &= \inf_{\lambda \in \Lambda_t^{*, \theta}(x, a)} \mathbb{E}_{X \sim \mathbb{P}_t^0(x, a, \cdot)} \left[ \sup_{y \in Y_{t, \theta, \lambda}^*(x, a; X)} D_\theta V_{t+1}^\theta(y)[r] \right]. \end{aligned}$$

# Robust actor–critic algorithm

**Approximators** updated backward in time,  $t = T - 1, \dots, 0$ :

$$V_{\psi,t}(x) \approx V_t^\theta(x), \quad U_{\xi,t}(x) \approx \nabla_\theta V_t^\theta(x).$$

**Per step:** draw  $a_t \sim \pi_t^\theta(x_t, \cdot)$  and  $X_{t+1}^{(i)} \sim \mathbb{P}_t^0(x_t, a_t, \cdot)$ ,  $i = 1, \dots, N$ .

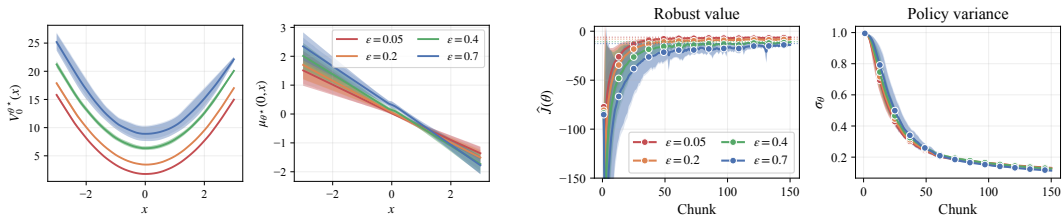
1. **Critic.** Solve the scalar dual on  $[0, \Lambda]$ ,  $\hat{\lambda}_t^* \in \arg \max_{\lambda \in [0, \Lambda]} \hat{F}_t^\theta(\lambda; x_t, a_t)$ , then regress  $V_{\psi,t}(x_t)$  on  $\hat{G}_t^\theta(x_t, a_t)$ .
2. **Inner selector.** For each  $i$ ,  $y^*(X_{t+1}^{(i)}) \in \arg \min_{y \in \mathcal{X}} \{f(t, x_t, a_t, y) + V_{\psi,t+1}(y) + \hat{\lambda}_t^* c(X_{t+1}^{(i)}, y)\}$ .
3. **Gradient critic.** Form  $\widehat{\nabla_\theta G_t^\theta}(x_t, a_t) = \frac{1}{N} \sum_{i=1}^N U_{\xi,t+1}(y^*(X_{t+1}^{(i)}))$ , then regress  $U_{\xi,t}(x_t)$  on  $\hat{z}_t = \hat{G}_t^\theta(x_t, a_t) \nabla_\theta \log \pi_t^\theta(x_t, a_t) + \widehat{\nabla_\theta G_t^\theta}(x_t, a_t)$ .
4. **Actor.**  $\theta \leftarrow \theta + \eta_\theta \cdot M^{-1} \sum_{j=1}^M U_{\xi,0}(X_0^{(j)})$ ,  $X_0^{(j)} \sim \mu_0$ .

*Remarks.* (i) No small- $\varepsilon$  restriction. (ii) Tabular case is exact; function approximation is a scalability device. (iii) Framework extends to any convex, weakly-compact ambiguity admitting strong duality.

## Numerical evidence — LQ control

**Setup.** Dynamics  $X_{t+1} = AX_t + Bu_t + \Xi w_t$ , quadratic costs  $x^\top Qx + u^\top Ru$ , terminal  $x^\top P_T x$ . Adversary perturbs the empirical noise distribution within a Wasserstein ball. Policy parametrized as a truncated Gaussian with neural-network mean  $\mu_\theta(t, x)$  and learned scalar variance  $\sigma_\theta^2$ .

**Recovery of the Riccati structure.** The learned  $V_0^{\theta^*}$  is quadratic in  $x$  and  $\mu_{\theta^*}(0, \cdot)$  is linear, matching the closed-form benchmark of Kim–Yang (2021):



Training is stable for all  $\varepsilon$ ; the policy variance  $\sigma_\theta$  collapses toward 0, so the stochastic policy becomes nearly deterministic, consistent with the deterministic Riccati optimizer  $K_t x + L_t$ .

**Take-away.** Wasserstein duality + directional Danskin  $\Rightarrow$  tractable robust actor–critic; recovers the Riccati optimum in LQ. [github.com/YadhHafsi/Policy-Gradient-Learning-for-DRO](https://github.com/YadhHafsi/Policy-Gradient-Learning-for-DRO)