

Constrained Policy Optimization: A Tale of Regularization and Optimism

Dongsheng Ding

dongshed@utk.edu



THE UNIVERSITY OF
TENNESSEE
KNOXVILLE

IFAC Workshop on Contextual Control; Aug. 23, 2026

Motivating application

■ ROBOTICS: ROBOT LOCOMOTION



Figure A1

maximize
policy

task progress

subject to

safety $\geq$ margin

CHALLENGE: constraint satisfaction

Context

■ SUCCESS STORIES OF RL

Go/Atari game, drone/car racing, LLMs, etc.

■ LESSONS LEARNED

- ★ importance of **policy optimization**

simple; scalable; model-free

- ★ **non-convex**; optimality w/ **one** performance metric

reward

- ★ **difficult** for **multiple** performance metrics

■ WHAT NOW ?

- ★ **applications**: robotics, healthcare, finance

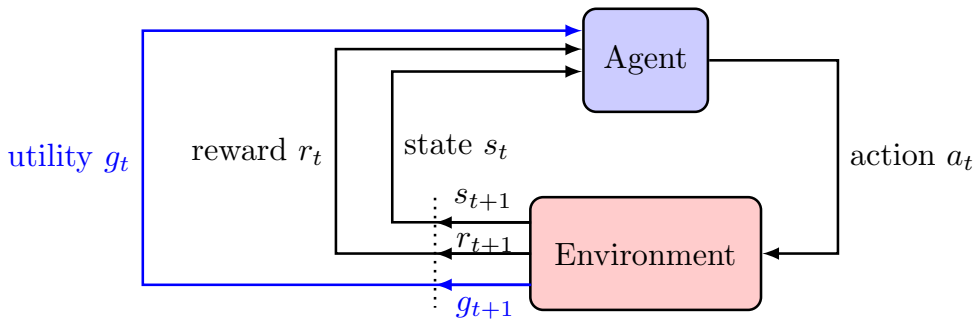
- ★ **policy optimization**: tremendous advances

OBJECTIVE

find a policy that
maximizes a performance metric
subject to a constraint on
another performance metric

Framework of RL

■ CONSTRAINED MDP



$\pi : \mathcal{S} \text{ (states)} \rightarrow \mathcal{A} \text{ (actions)}$ – a policy

$V_r^\pi(\rho) := \mathbb{E}[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 \sim \rho]$ – reward value function

$V_g^\pi(\rho) := \mathbb{E}[\sum_{t=0}^{\infty} \gamma^t g(s_t, a_t) \mid s_0 \sim \rho]$ – utility value function

Constrained policy optimization

maximize $V_r^\pi(\rho)$ $\longrightarrow$ **standard objective**

subject to $V_g^\pi(\rho) \geq 0$

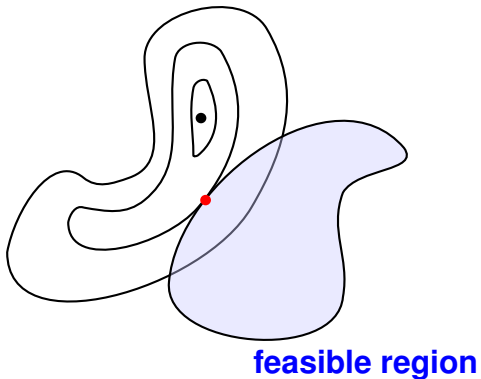


policy constraint

- ★ limit the policy space to an inequality constraint

■ 2-D LEVEL NONCONVEX CURVES

optimal, but in**feasible**



optimality w/ **feasibility**

$\Rightarrow$ init. dependent & stochastic

Lagrangian approach

Lagrangian: $L(\pi, \lambda) = V_r^\pi(\rho) + \lambda V_g^\pi(\rho)$

Lagrange multiplier: $\lambda (\geq 0)$

composite reward: $r + \lambda g$

■ POLICY OPTIMIZATION

$$\underset{\pi}{\text{maximize}} \quad V_{r+\lambda g}^\pi(\rho)$$

multiple solutions, but **infeasible**

Primal-dual method

primal problem: $\max_{\pi} \min_{\lambda} V_{r+\lambda g}^{\pi}(\rho)$

dual problem: $\min_{\lambda} \max_{\pi} V_{r+\lambda g}^{\pi}(\rho)$

■ MIN-MAX POLICY OPTIMIZATION

strong duality $\implies$ **exchange of min and max**

$$\min_{\lambda} \max_{\pi} V_{r+\lambda g}^{\pi}(\rho) = \max_{\pi} \min_{\lambda} V_{r+\lambda g}^{\pi}(\rho)$$

OBJECTIVE: find a minimax point (π^*, λ^*)

■ SIMULTANEOUS ASCENT-DESCENT

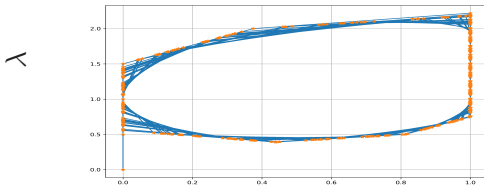
policy ascent direction: $\nabla_{\pi} V_{r+\lambda g}^{\pi}(\rho)$

dual descent direction: $\nabla_{\lambda} V_{r+\lambda g}^{\pi}(\rho)$

$$\pi^+ \leftarrow \pi + \eta \nabla_{\pi} V_{r+\lambda g}^{\pi}(\rho)$$

$$\lambda^+ \leftarrow \lambda - \eta \nabla_{\lambda} V_{r+\lambda g}^{\pi}(\rho)$$

single-time-scale w/ stepsize η



$$\pi \left(\pi^* = 0.66 \right)$$

CHALLENGE

single-time-scale primal-dual methods
in face of **nonconvex** Lagrangian

Glimpse of our efforts

policy convergence w/ **(sub) linear** error rate

■ REGULARIZED POLICY GRADIENT PRIMAL-DUAL METHOD

★ **tabular**

dimension-free

★ function approximation

up to approx. error

■ OPTIMISTIC POLICY GRADIENT PRIMAL-DUAL METHOD

★ **tabular**

problem-dependent

DWZR, NeurIPS 2023

DHR, AISTATS 2024

R*D*MR, AAAI 2025

Regularized method

Regularized Lagrangian approach

entropy-like term: $\mathcal{H}(\pi) = \mathbb{E} \left[\sum_{t=0}^{\infty} -\gamma^t \log \pi(a_t | s_t) \right]$

■ REGULARIZED LAGRANGIAN

$$L_{\tau}(\pi, \lambda) = V_{r+\lambda g}^{\pi}(\rho) + \tau \left(\mathcal{H}(\pi) + \frac{1}{2} \lambda^2 \right)$$



add stabilizing regularization

OBJECTIVE: find the regularized minimax point $(\pi_{\tau}^*, \lambda_{\tau}^*)$

★ **τ -near minimax point of** $V_{r+\lambda g}^{\pi}(\rho)$

Two pillars

■ Q-VALUE FUNCTION

$$Q_r^\pi(s, a) = \mathbb{E} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \mid s_0 = s, a_0 = a \right]$$

■ STATE VISITATION DISTRIBUTION

$$d_{s_0}^\pi(s) = (1 - \gamma) \sum_{t=0}^{\infty} \gamma^t P^\pi(s_t = s \mid s_0)$$

★ expectation $d_\rho^\pi(s) = \mathbb{E}_{s_0 \sim \rho} [d_{s_0}^\pi(s)]$

Regularized policy gradient primal-dual method

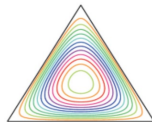
policy ascent direction: $Q_{L_\tau}^\pi(s, a) := Q_{r + \lambda g - \tau \log \pi}(s, a)$

dual descent direction: $V_g^\pi(\rho) + \tau \lambda$

■ SINGLE-TIME-SCALE PRIMAL-DUAL UPDATE

$$\begin{aligned}\pi^+(\cdot | s) &= \operatorname{argmax}_{\pi' \in \Delta_{\epsilon_0}} \langle \pi'(\cdot | s), Q_{L_\tau}^\pi(s, \cdot) \rangle - \frac{1}{\eta} \text{KL}_s(\pi', \pi) \\ \lambda^+ &= \operatorname{argmin}_{\lambda' \in \Lambda} \lambda' (V_g^\pi(\rho) + \tau \lambda) + \frac{1}{2\eta} (\lambda' - \lambda)^2\end{aligned}$$

★ restricted policy set $\Delta_{\epsilon_0} = \{p_a \geq \epsilon_0, \forall a\}$



Non-asymptotic last-iterate performance

distance of (π_t, λ_t) to $(\pi_\tau^*, \lambda_\tau^*)$: $\mathbb{E}_{s \sim d_\rho^{\pi_\tau^*}} [\text{KL}_s(\pi_\tau^*, \pi_t)] + \frac{1}{2}(\lambda_t - \lambda_\tau^*)^2$

Theorem (informal)

★ **distance of (π_t, λ_t) to $(\pi_\tau^*, \lambda_\tau^*)$ is bounded by**

$$e^{-\eta \tau t} + \frac{\eta}{\tau}$$

exponential decay up to a ratio

★ ϵ -near regularized minimax point requires

$$\eta = \epsilon \tau \quad \text{and} \quad t = \frac{1}{\epsilon \tau^2} \quad \text{sublinear rate}$$

optimality gap: $V_r^*(\rho) - V_r^{\pi_t}(\rho)$

feasibility gap: $0 - V_g^{\pi_t}(\rho)$

Implication (informal)

★ ϵ -optimality gap & ϵ -feasibility gap require

$\frac{1}{\epsilon^6}$ iterations

sublinear rate

stepsize $\eta = \epsilon^4$

regularization $\tau = \epsilon^2$

Distance of primal-dual iterates to $(\pi_\tau^*, \lambda_\tau^*)$

Step #1: performance difference & MWU

$$L_\tau(\pi_\tau^*, \lambda_t) - L_\tau(\pi_t, \lambda_t)$$

$$\lesssim \mathbb{E}_{s \sim d_\rho^{\pi_\tau^*}} [\langle \pi_\tau^*(\cdot | s) - \pi_t(\cdot | s), Q_{r+\lambda_t g - \tau \pi_t}^{\pi_t}(s, \cdot) \rangle] - \tau \mathbb{E}_{s \sim d_\rho^{\pi_\tau^*}} [\text{KL}_t(s)]$$

$$\lesssim \frac{1}{\eta} \mathbb{E}_{s \sim d_\rho^{\pi_\tau^*}} [\text{KL}_t(s) - \text{KL}_{t+1}(s)] + \eta(C_\tau)^2 - \tau \mathbb{E}_{s \sim d_\rho^{\pi_\tau^*}} [\text{KL}_t(s)]$$

$$= \frac{1}{\eta} [(1 - \eta\tau) \text{KL}_t(\rho) - \text{KL}_{t+1}(\rho)] + \eta(C_\tau)^2$$

$$\text{KL}_t(s) = \text{KL}(\pi_\tau^*(\cdot | s), \pi_t(\cdot | s))$$

$$\text{KL}_t(\rho) = \mathbb{E}_{s \sim d_\rho^{\pi_\tau^*}} [\text{KL}_t(s)]$$

$$\boxed{L_\tau(\pi_t, \lambda_t) - L_\tau(\pi_t, \lambda_\tau^*)}$$

$$\lesssim \frac{1}{\eta} \left[(1 - \eta\tau)(\lambda_\tau^* - \lambda_t) - (\lambda_\tau^* - \lambda_{t+1}) \right] + \eta(C'_\tau)^2$$

Step #2: saddle point property & geometric series

$$\cancel{\eta (L_\tau(\pi_\tau^*, \lambda_t) - L_\tau(\pi_t, \lambda_\tau^*))} + \boxed{\text{KL}_{t+1}(\rho) + \frac{1}{2}(\lambda_\tau^* - \lambda_{t+1})^2}$$

$$\lesssim \boxed{(1 - \eta\tau) \left(\text{KL}_t(\rho) + \frac{1}{2}(\lambda_\tau^* - \lambda_t)^2 \right)} + \eta^2 \max \left((C_\tau)^2, (C'_\tau)^2 \right)$$

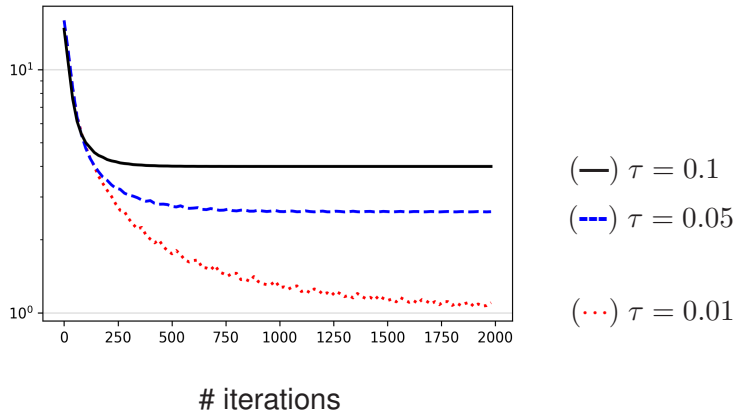
$\vdots$

$$\lesssim (1 - \eta\tau)^t \left(\text{KL}_1(\rho) + \frac{1}{2}(\lambda_\tau^* - \lambda_1)^2 \right)$$

$$+ \left(\eta^2 \left(1 + (1 - \eta\tau) + (1 - \eta\tau)^2 + \dots \right) \right) \max \left((C_\tau)^2, (C'_\tau)^2 \right)$$

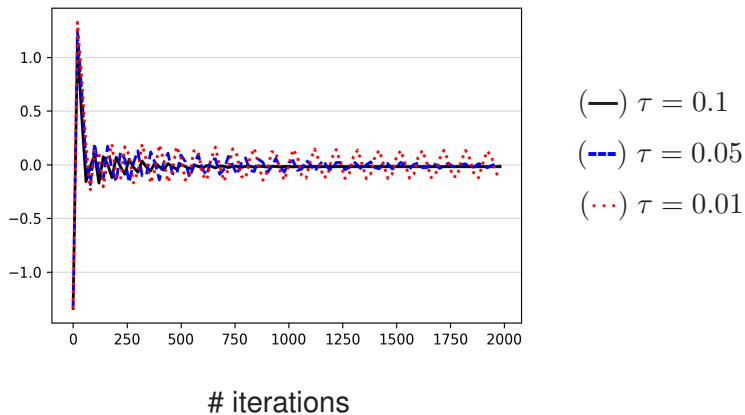
Sublinear convergence

$$\sum_s \|\pi_t(\cdot | s) - \pi^*(\cdot | s)\|^2$$



smaller regularization $\tau \longrightarrow$ **better** accuracy

$$V_g^{\pi_t}(\rho)$$

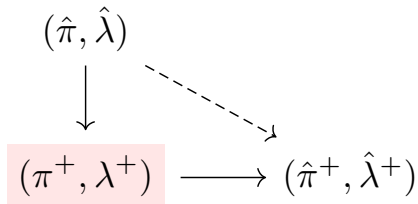


smaller regularization $\tau \longrightarrow$ **larger** oscillation

■ ACCURACY VS OSCILLATION TRADE-OFF

Optimistic method

Optimistic policy gradient primal-dual method



policy ascent direction: $Q_{r+\hat{\lambda}g}^{\hat{\pi}}(s, a)$

dual descent direction: $V_g^{\hat{\pi}}(\rho)$

policy ascent direction: $Q_{r+\lambda^+g}^{\pi^+}(s, a)$

dual descent direction: $V_g^{\pi^+}(\rho)$

■ SINGLE-TIME-SCALE PRIMAL-DUAL UPDATE

prediction step

$$\pi^+(\cdot | s) = \operatorname{argmax}_{\pi' \in \Delta} \langle \pi'(\cdot | s), Q_{r+\lambda g}^\pi(s, \cdot) \rangle - \frac{1}{2\eta} \|\pi' - \hat{\pi}\|_s^2$$

$$\lambda^+ = \operatorname{argmin}_{\lambda' \in \Lambda} \lambda' V_g^\pi(\rho) + \frac{1}{2\eta} (\lambda' - \hat{\lambda})^2$$

real step

$$\hat{\pi}^+(\cdot | s) = \operatorname{argmax}_{\pi' \in \Delta} \langle \pi'(\cdot | s), Q_{r+\lambda^+ g}^{\pi^+}(s, \cdot) \rangle - \frac{1}{2\eta} \|\pi' - \hat{\pi}\|_s^2$$

$$\hat{\lambda}^+ = \operatorname{argmin}_{\lambda' \in \Lambda} \lambda' V_g^{\pi^+}(\rho) + \frac{1}{2\eta} (\lambda' - \hat{\lambda})^2$$

Non-asymptotic last-iterate performance

distance of $(\hat{\pi}_t, \hat{\lambda}_t)$ to $\Pi^* \times \Lambda^*$: $\mathbb{E}_{s \sim d_{\rho}^{\pi^*}} [\text{Dist}_s(\hat{\pi}_t, \Pi^*)] + \text{Dist}(\hat{\lambda}_t, \Lambda^*)$

Theorem (informal)

★ **distance of $(\hat{\pi}_t, \hat{\lambda}_t)$ to $\Pi^* \times \Lambda^*$ is bounded by**

$$\left(\frac{1}{1 + C} \right)^t$$

exponential decay to zero

problem-dependent constants C, η

optimality gap: $V_r^*(\rho) - V_r^{\pi_t}(\rho)$

feasibility gap: $0 - V_g^{\pi_t}(\rho)$

Implication (informal)

★ **ϵ -optimality gap & ϵ -feasibility gap** require

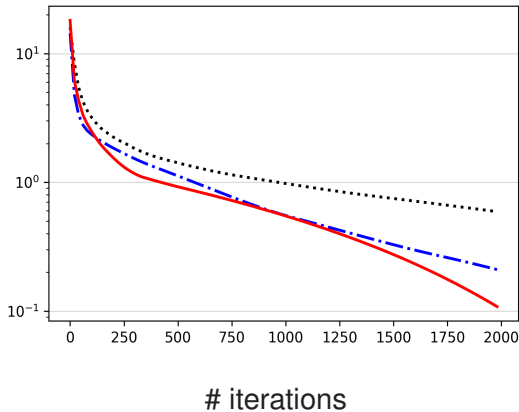
$\log^2 \frac{1}{\epsilon}$ iterations

linear rate

problem-dependent

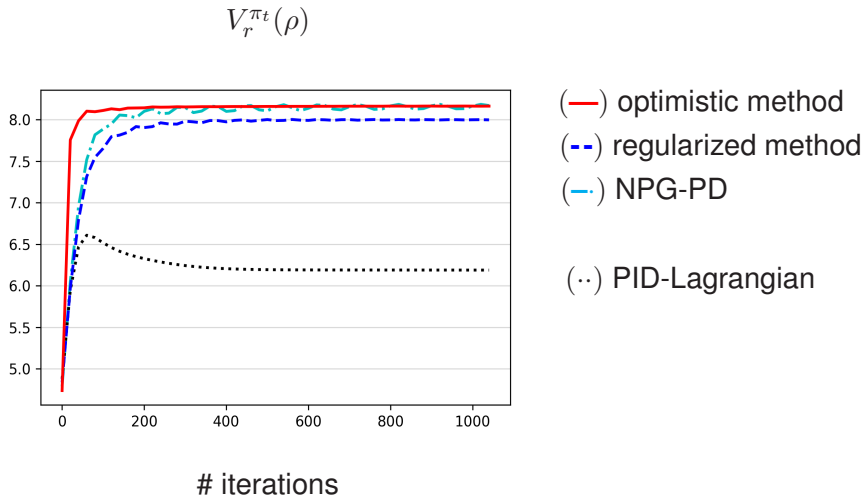
Linear convergence

$$\sum_s \|\hat{\pi}_t(\cdot | s) - \pi^*(\cdot | s)\|^2$$



linear rate holds for a range of stepsizes

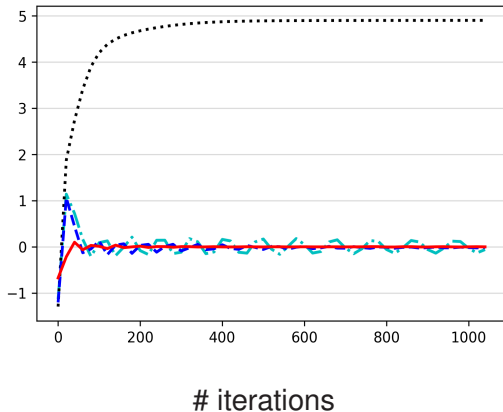
Comparison of primal-dual methods



optimistic method $\approx$ NPG-PD $\geq$ regularized method $\geq$ PID-Lagrangian

Comparison of primal-dual methods

$$V_g^{\pi^t}(\rho)$$



(..) PID-Lagrangian

(—) optimistic method

(- -) regularized method

(- ·) NPG-PD

optimistic method $\approx$ regularized method $\leq$ NPG-PD $\leq$ PID-Lagrangian

Summary

■ SINGLE-TIME-SCALE PRIMAL-DUAL METHODS

- ★ regularized policy gradient primal-dual method
- ★ optimistic policy gradient primal-dual method
- ★ non-asymptotic last-iterate policy convergence

■ ON-GOING EFFORTS

- ★ constrained policy optimization w/o exploration
- ★ more practical considerations

Thank you for your attention.

Intuition of exponential decay

Distance of primal-dual iterates to $\Pi^* \times \Lambda^*$

Step #1: monotonic decreasing distance

$$\begin{aligned}\Theta_t &= \frac{1}{2} \mathbb{E}_{s \sim d_{\rho}^{\pi^*}} \|\hat{\pi}_t(\cdot | s) - \mathcal{P}_{\Pi^*}(\hat{\pi}_t(\cdot | s))\|^2 + \frac{1}{2} \left(\hat{\lambda}_t - \mathcal{P}_{\Lambda^*}(\hat{\lambda}_t) \right)^2 \\ &\quad + \frac{1}{16} \mathbb{E}_{s \sim d_{\rho}^{\pi^*}} \|\hat{\pi}_t(\cdot | s) - \pi_{t-1}(\cdot | s)\|^2 + \frac{1}{2} \left(\hat{\lambda}_t - \lambda_{t-1} \right)^2 \\ &\geq \Theta_{t+1} + \frac{7}{16} \underbrace{\zeta_t}_{\text{blue}} \\ &= \frac{1}{2} \mathbb{E}_{s \sim d_{\rho}^{\pi^*}} \|\hat{\pi}_{t+1}(\cdot | s) - \pi_t(\cdot | s)\|^2 + \frac{1}{2} \left(\hat{\lambda}_{t+1} - \lambda_t \right)^2 \\ &\quad + \frac{1}{2} \mathbb{E}_{s \sim d_{\rho}^{\pi^*}} \|\hat{\pi}_t(\cdot | s) - \pi_t(\cdot | s)\|^2 + \frac{1}{2} \left(\hat{\lambda}_t - \lambda_t \right)^2\end{aligned}$$

Step #2: lower bound of ζ_t

$$\begin{aligned}
 \zeta_t &\gtrsim \frac{1}{4} \mathbb{E}_{s \sim d_{\rho}^{\pi^*}} \|\hat{\pi}_{t+1}(\cdot | s) - \pi_t(\cdot | s)\|^2 + \frac{1}{4} \left(\hat{\lambda}_{t+1} - \lambda_t \right)^2 \\
 &\quad + \eta^2 \underbrace{\sup_{\pi, \lambda} \frac{\left[L(\pi, \hat{\lambda}_{t+1}) - L(\hat{\pi}_{t+1}, \lambda) \right]_+^2}{\max_s \left(\|\pi(\cdot | s) - \hat{\pi}_{t+1}(\cdot | s)\| + |\lambda - \hat{\lambda}_{t+1}| \right)^2}}_{\text{problem-dependent constant}} \\
 &\gtrsim \mathbb{E}_{s \sim d_{\rho}^{\pi^*}} \|\hat{\pi}_{t+1}(\cdot | s) - \mathcal{P}_{\Pi^*}(\hat{\pi}_{t+1}(\cdot | s))\|^2 + \left(\hat{\lambda}_{t+1} - \mathcal{P}_{\Lambda^*}(\hat{\lambda}_{t+1}) \right)^2 \\
 &\gtrsim \eta^2 \Theta_{t+1}
 \end{aligned}$$

$$\implies \boxed{\Theta_{t+1} \lesssim \Theta_t - \eta^2 \Theta_{t+1}}$$