

RL with continuous action spaces

Linear MDP

Aarti Singh

Machine Learning 10-734

Oct 30, 2025

Slides courtesy: Yuejie Chi, Wen Sun



MACHINE LEARNING DEPARTMENT

Carnegie Mellon.
School of Computer Science

Continuous action spaces

Bandits:

Reward is linear, Lipschitz, GP, NN, ...

e.g. $r^*(x) = x^\top \theta^*$ x, θ^* are d-dimensional

Why?

Optimal action a^* depends on reward r^*

MDP:

Optimal policy π^* depends on Q^*

Attempt 1: Value function Q^* is linear

e.g. $Q^*(s,a) = \phi(s,a)^\top \theta^*$ ϕ, θ^* are d-dimensional

feature representation of (state, action) pair

Linear Q doesn't suffice

- No, regret bounds:

	Linear Bandits	Q^* Linear
Upper bound	$\tilde{O}(d\sqrt{T})$ LinUCB	$O(e^H)$
Lower bound	$\Omega(d\sqrt{T})$	$\Omega(\text{poly}(H))$

- Need additional assumptions!

Linear MDP

- Attempt 2: Use a parametric function to approximate $r(s, a)$, $P(s'|s, a)$

- E.g. Linear $r_h(s, a) = w_h^T \phi(s, a)$

$$\text{Low-rank } P_h(s'|s, a) = \mu_h(s')^T \phi(s, a)$$

- Also called low-rank MDPs

$$\begin{array}{ccc} \boxed{\begin{array}{c} S \times S \\ P(s'|s, \pi(s)) \end{array}} & = & \boxed{\begin{array}{c} S \times d \\ \Psi^T \end{array}} \quad \boxed{\begin{array}{c} d \times S \\ \Phi^\pi \end{array}} \end{array}$$

d-rank

- For all π ,

parameters

Linear MDP

- Use a parametric function to approximate $r(s, a), P(s'|s, a)$
 - E.g. Linear $r_h(s, a) = w_h^T \phi(s, a)$

$$\text{Low-rank } P_h(s'|s, a) = \mu_h(s')^T \phi(s, a)$$

- Learn the parameters , *value iteration* *policy iteration*
- Training data $(\mathbf{s}_t, a_t, r_t, \mathbf{s}_{t+1})$ gathered online by the agent/learning algorithm

Fundamental property of Linear MDP

For any value function V , Bellman update $T_h V$ is linear

where Bellman operator T_h

Bellman update

$$T_h V(s) = r_h(s, \pi(s)) + \sum_{s'} P(s' | s, \pi(s)) V^\pi(s')$$

Proof: $r_h(s, a) = w_h^T \phi(s, a)$

$$\sum_{s'} P(s' | s, a) V^\pi(s') = \sum_{s'} \underbrace{\mu_h(s')^T \phi(s, a)}_{\text{equiv. parameter for } V \text{ function}} V^\pi(s')$$

$$= \sum_{s'} \underbrace{\mu_h(s')^T V^\pi(s')}_{\text{equiv. parameter for } V \text{ function}} \cdot \phi(s, a)$$

equiv. parameter for V function

Implies Q^* is linear! But stronger

$$V^* = T_h V^*, \quad Q^* = T Q^*$$



LSVI-UCB: Least Square Value Iteration with UCB

$s_h^i, a_h^i, r_h^i, s_{h+1}^i$

Value iteration at episode n using $\{s_h^i, a_h^i, r_h^i, s_{h+1}^i\}_{h=1, i=1}^{H-1, n-1}$

$$\widehat{V}_H^n(s) = 0, \forall s$$

For $h = H-1, H-2, \dots, 1$

$$\theta_h^n \leftarrow \underset{\theta}{\operatorname{argmin}} \sum_{i=1}^{n-1} \left(\underbrace{\langle \phi(s_h^i, a_h^i), \theta \rangle}_{\langle X_i, \theta \rangle} - \underbrace{r_h^i - \widehat{V}_{h+1}^n(s_{h+1}^i)}_{Y_i} \right)^2 + \lambda \|\theta\|_2^2$$

$$\widehat{Q}_h^n(s, a) = \min \left\{ \underbrace{b_h^n(s, a)}_{\text{lower bound}} + \langle \phi(s, a), \theta_h^n \rangle, H \right\}, \forall s, a$$

$\widehat{Q} > 0^*$
upper confidence bound

$$\widehat{V}_h^n(s) = \max_a \widehat{Q}_h^n(s, a), \quad \pi_h^n(s) = \arg \max_a \widehat{Q}_h^n(s, a), \forall s$$

LSVI-UCB: Least Square Value Iteration with UCB

At each episode n

For $h = H, H-1, \dots, 1$

Estimate using data from past episodes

$$\widehat{Q}_h^n(s, a) = \min \left\{ b_h^n(s, a) + \langle \phi(s, a), \theta_h^n \rangle, H \right\}, \forall s, a$$

Plan

$$\pi_h^n(s) = \arg \max_a \widehat{Q}_h^n(s, a), \forall s$$

Execute

Roll out π_h^n to collect trajectory and add to data

LSVI-UCB: Least Square Value Iteration with UCB

$$\widehat{Q}_h^n(s, a) = \min \left\{ b_h^n(s, a) + \langle \phi(s, a), \theta_h^n \rangle, H \right\}, \forall s, a$$

$$\theta_h^n \leftarrow \operatorname{argmin}_{\theta} \sum_{i=1}^{n-1} \left(\underbrace{\langle \phi(s_h^i, a_h^i), \theta \rangle}_{\langle x_i, \theta \rangle} - r_h^i - \underbrace{\widehat{V}_{h+1}^n(s_{h+1}^i)}_{y_i} \right)^2 + \lambda \|\theta\|_2^2$$

Ridge

This implies $\theta_h^n = \Lambda_h^{n-1} \sum_{i=1}^{n-1} \phi_h^i (r_h^i + \widehat{V}_{h+1}^n(s_{h+1}^i))$

$\left(\sum_{i=1}^n x_i x_i^T + \lambda I \right)^{-1} \sum_{i=1}^n x_i y_i$
 Σ

where $\phi_h^i = \phi(s_h^i, a_h^i)$ and $\Lambda_h^n = \sum_{i=1}^{n-1} \phi_h^i \phi_h^{iT} + \lambda I$

and bonus $b_h^n(s, a) = \|\phi(s, a)\|_{\Lambda_h^{n-1}} \beta$ (derived later)

$\|x\|_{\Sigma^{-1}} \beta$

Bandits: If $H = 1$, recover LinUCB for Linear bandits

LSVI-UCB Regret

- With probability at least $1 - \delta$, LSVI-UCB has regret

$$\text{Regret}, \sum_{n=1}^N [\underline{V}^* - \underline{V}^{\pi_n}] = \tilde{O}(H^2 \sqrt{d^3 T})$$

Key proof idea: Combine proofs of LinUCB and UCB-VI

↓
bandits

↓
Tabular

LSVI-UCB Regret

- With probability at least $1 - \delta$, LSVI-UCB has regret

$$\text{Regret}, \sum_{n=1}^N [\underline{V^* - V^{\pi_n}}] = \tilde{O}(H^2 \sqrt{d^3 T})$$

$$\begin{aligned} V^* - V^{\pi_n} &\leq \hat{V}^n - V^{\pi_n} \\ &\leq \text{sim lemma} \end{aligned}$$

Step 1: Design bonus to guarantee Optimism i.e. $V_h^*(s) \leq \hat{V}_h^n(s) \forall h, s$
 implies Bonus $b_h^n(s, a) = \text{Transition error } \underline{e_h^n(s, a)}$

Step 2: Characterize transition error

Step 3: Apply simulation lemma

$$\hat{V}_0^n(s_0) - V_0^{\pi_n}(s_0) \leq \sum_{h=0}^{H-1} \mathbb{E}_{s,a \sim d_h^{\pi_n}} \left[b_h^n(s, a) + \underbrace{(\hat{P}_h^n(\cdot | s, a) - P_h(\cdot | s, a)) \cdot \hat{V}_{h+1}^n}_{e_h^n(s, a)} \right]$$

Handwritten notes:
 $\hat{V}_0^n$ (under $\hat{V}_0^n$)
 $V_0^{\pi_n}$ (under $V_0^{\pi_n}$)
 $e_h^n(s, a)$ (under $e_h^n(s, a)$)
 Transition error (under $e_h^n(s, a)$)

LSVI-UCB Regret

- With probability at least $1 - \delta$, LSVI-UCB has regret

$$\text{Regret}, \sum_{n=1}^N [V^* - V^{\pi_n}] = \tilde{O}(H^2 \sqrt{d^3 T})$$

Step 1: Design bonus to guarantee Optimism i.e. $V_h^*(s) \leq \hat{V}_h^n(s) \forall h, s$

For any policy π , consider

$$\begin{aligned} \hat{Q}_h^n(s, a) - Q_h^\pi(s, a) & \quad \text{def of } \hat{Q} \\ &= b_h^n(s, a) + \langle \phi(s, a), \theta_h^n \rangle - r_h(s, a) - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] \\ &= \underbrace{b_h^n(s, a) + \langle \phi(s, a), w_h + \theta_h^{\hat{V}} \rangle}_{\text{Bellman update for true value}} + \underbrace{\langle \phi(s, a), \Lambda_h^{n-1} \sum_{i=1}^{N-1} \phi_h^i \eta_h^i \rangle}_{\text{transition error}} \\ & \quad - r_h(s, a) - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] \end{aligned}$$

$$\begin{aligned} \theta_h^n &= \Lambda_h^{n-1} \sum_{i=1}^{n-1} \phi_h^i (r_h^i + \hat{V}_{h+1}^i(s_{h+1}^i)) \\ &= \Lambda_h^{n-1} \sum_{i=1}^{n-1} \phi_h^i (r_h^i + \mathbb{E}_{s' \sim P(\cdot | s_h^i, a_h^i)} [\hat{V}_{h+1}^i(s')] + \eta_h^i) \quad \text{where } \eta_h^i \text{ is transition error} \\ &= \Lambda_h^{n-1} \sum_{i=1}^{n-1} \phi_h^i (\phi_h^i w_h + \phi_h^i \theta_h^{\hat{V}} + \eta_h^i) \quad \text{Bellman update is linear in linear MDP} \end{aligned}$$

$$= \cancel{\Lambda_h^{n-1}} \cancel{\sum_{i=1}^{n-1}} (\underbrace{w_h + \theta_h^{\hat{V}}}_{\text{Bellman update for true value}}) + \Lambda_h^{n-1} \sum_{i=1}^{n-1} \phi_h^i \eta_h^i$$

LSVI-UCB Regret

- With probability at least $1 - \delta$, LSVI-UCB has regret

$$\text{Regret}, \sum_{n=1}^N [V^* - V^{\pi_n}] = \tilde{O}(H^2 \sqrt{d^3 T})$$

Step 1: Design bonus to guarantee Optimism i.e. $V_h^*(s) \leq \hat{V}_h^n(s) \forall h, s$

For any policy π , consider

$$\begin{aligned}
 & \hat{Q}_h^n(s, a) - Q_h^\pi(s, a) \\
 &= b_h^n(s, a) + \langle \phi(s, a), \theta_h^n \rangle - r_h(s, a) - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] \\
 &= b_h^n(s, a) + \langle \phi(s, a), w_h + \theta_h^{\hat{V}} \rangle + \langle \phi(s, a), \Lambda_h^{n-1} \sum_{i=1}^{N-1} \phi_h^i \eta_h^i \rangle - r_h(s, a) \\
 &\quad - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] \\
 &= b_h^n(s, a) + \langle \phi(s, a), \theta_h^{\hat{V}} \rangle + \langle \phi(s, a), \Lambda_h^{n-1} \sum_{i=1}^{N-1} \phi_h^i \eta_h^i \rangle \\
 &\quad - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] \\
 &= b_h^n(s, a) + \mathbb{E}_{s' \sim P(\cdot | s, a)} [\hat{V}_{h+1}^i(s')] + \langle \phi(s, a), \Lambda_h^{n-1} \sum_{i=1}^{N-1} \phi_h^i \eta_h^i \rangle \\
 &\quad - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] \\
 &= b_h^n(s, a) + \mathbb{E}_{s' \sim P(\cdot | s, a)} [\hat{V}_{h+1}^i(s')] - \mathbb{E}_{s' \sim P(\cdot | s, a)} [V_{h+1}^\pi(s')] + e_h^n(s, a) \\
 &\geq 0. \quad \text{by induction for optimism and choice of bonus}
 \end{aligned}$$

$b - e \geq 0$
choose bonus

$$\text{if } \hat{V}_{h+1} \geq V_{h+1} \Rightarrow \hat{Q}_h \geq Q_h \Rightarrow \hat{V}_h \geq V_h$$

LSVI-UCB Regret

- With probability at least $1 - \delta$, LSVI-UCB has regret

$$\text{Regret}, \sum_{n=1}^N [V^* - V^{\pi_n}] = \tilde{O}(H^2 \sqrt{d^3 T})$$

Step 2: Characterize transition error

$$\text{w.p. at least } 1-\delta, |e_h^n(s, a)| = |\phi^T(s, a) \Lambda_h^{n-1} \sum_{i=1}^{n-1} \phi_h^i \eta_h^i|$$

$$\leq \|\phi\|_{\Lambda_h^{n-1}} \left\| \sum_{i=1}^{n-1} \phi_h^i \eta_h^i \right\|_{\Lambda_h^{n-1}}$$

$$\leq \|\phi\|_{\Lambda_h^{n-1}} \beta \quad \leftarrow \text{bonus}$$

where $\left\| \sum_{i=1}^{n-1} \phi_h^i \eta_h^i \right\|_{\Lambda_h^{n-1}} \leq \beta = \tilde{O}(dH)$ from LinUCB proof

$$\underbrace{\phi^T \sum_x^{-1} \sum_{i=1}^n x_i^T \varepsilon_i}_{\text{LinUCB}}$$

LinUCB

LSVI-UCB Regret

- With probability at least $1 - \delta$, LSVI-UCB has regret

$$\text{Regret}, \sum_{n=1}^N [V^* - V^{\pi_n}] = \tilde{O}(H^2 \sqrt{d^3 N}) \quad \checkmark \quad \text{poly}(H)$$

Step 3: Apply simulation lemma

$$\widehat{V}_0^n(s_0) - V_0^{\pi^n}(s_0) \leq \sum_{h=0}^{H-1} \mathbb{E}_{s,a \sim d_h^{\pi^n}} \left[b_h^n(s, a) + \overbrace{(\widehat{P}_h^n(\cdot | s, a) - P_h(\cdot | s, a)) \cdot \widehat{V}_{h+1}^n}^{e_h^n(s, a)} \right]$$

$$\sum_{n=1}^N [V^* - V^{\pi_n}] \leq \sum_{n=1}^N [V^* - V^{\pi_n}] \leq \sum_{n=1}^N \sum_{h=0}^{H-1} \mathbb{E}[b_h^n(s_h^n, a_h^n) + e_h^n(s_h^n, a_h^n)]$$

optimism

simulation

$2\beta \|\phi\|_{\Lambda^{-1}}$

$\beta = O(Hd)$

$$= 2\beta \sum_{h=0}^{H-1} \sqrt{N} \sqrt{\sum_{n=1}^N \|\phi_h^n\|_{\Lambda_h^{n-1}}^2} = \tilde{O}(\beta H \sqrt{dN}) = \tilde{O}(H^2 d \sqrt{dN}) \quad \checkmark$$

Cauchy-Schwarz

$O(d)$ using elliptical potential lemma from LinUCB

Fundamental property of Linear MDP

For any value function V , Bellman update $T_h V$ is linear
where Bellman operator

$$T_h V(s) = r_h(s, \pi(s)) + \sum_{s'} P(s' | s, \pi(s)) V^\pi(s') \quad \text{||}$$

Proof: $r_h(s, a) = \underline{w_h^T \phi(s, a)}$

$$\begin{aligned} \sum_{s'} P(s' | s, a) V^\pi(s') &= \sum_{s'} \mu_h(s')^T \phi(s, a) V^\pi(s') \\ &= \sum_{s'} \underbrace{\mu_h(s')^T V^\pi(s')}_{\theta_h^V} \cdot \phi(s, a) \end{aligned}$$

Implies Q^* is linear! But stronger

LSVI-UCB: Least Square Value Iteration with UCB

Value iteration at episode n using $\{s_{hi}, a_{hi}, r_{hi}, s_{h+1}^i\}_{h=1, i=1}^{H-1, n-1}$

$$\widehat{V}_H^n(s) = 0, \forall s$$

For $h = H-1, H-2, \dots, 1$

$$\theta_h^n \leftarrow \underset{\theta}{\operatorname{argmin}} \sum_{i=1}^{n-1} \left(\langle \phi(s_h^i, a_h^i), \theta \rangle - r_h^i - \widehat{V}_{h+1}^n(s_{h+1}^i) \right)^2 + \lambda \|\theta\|_2^2$$

LS uses linear Bellman update property

$$\widehat{Q}_h^n(s, a) = \min \left\{ b_h^n(s, a) + \langle \phi(s, a), \theta_h^n \rangle, H \right\}, \forall s, a$$

$$\widehat{V}_h^n(s) = \max_a \widehat{Q}_h^n(s, a), \quad \pi_h^n(s) = \arg \max_a \widehat{Q}_h^n(s, a), \forall s$$

RL with continuous action spaces

- Approach 3:

$\neq Q^*$ is linear

Linear Bellman completeness: For any linear Q function, its Bellman update is also linear

Given feature ϕ , take any linear function $w^\top \phi(s, a)$:

$$\forall h, \exists \theta \in \mathbb{R}^d, s.t., \theta^\top \phi(s, a) = r(s, a) + \underbrace{\mathbb{E}_{s' \sim P_h(s, a)} \max_{a'} w^\top \phi(s', a')}_{\text{Bellman update}}, \forall s, a$$

stronger than
assuming Q^* linear

but

weaker than r is linear
& P is low rank

RL with continuous action spaces

General function approximation

- Use a parametric function, $\underline{Q(s, a; \Theta)}$, to approximate $\underline{Q^*(s, a)}$
 - E.g. $Q(s, a; \Theta)$ = Neural network
Linear $\theta^T \phi(s, a)$ where $\phi(s, a)$ - features
- Learn the parameters *, policy value iteration*
- Training data $(\mathbf{s}_t, a_t, r_t, \mathbf{s}_{t+1})$ gathered online by the agent/learning algorithm

RL with continuous action spaces

General function approximation

- Use a parametric function, $Q(s, a; \Theta)$, to approximate $Q^*(s, a)$
 - E.g. $Q(s, a; \Theta)$ = Neural network
Linear $\theta^T \phi(s, a)$ where $\phi(s, a)$ - features

Q^* is realizable not enough!

Need **Bellman completeness**: For any Q function in $\mathcal{F}$, its Bellman update is also in $\mathcal{F}$