

Standard MDP

State :	S
Actions :	A
Reward :	$r(s,a)$
Transition :	$P(S' S,a)$
Discount :	γ
optimal policy :	$\pi^* = \arg \max_{\pi} \sum_a P^{\pi} E_{S \sim \pi} [r(S,a)]$

Maximum Entropy MDP

State :	S
Actions :	A
Reward :	$r(s,a) + \alpha \mathcal{H}(\pi(s))$
Transition :	$P(S' S,a)$
Discount :	γ
optimal policy :	$\pi_{SE}^* = \arg \max_{\pi} \sum_a P^{\pi} E_{S \sim \pi} [r(S,a) + \alpha \mathcal{H}(\pi(s))]$

Definition:

optimal policy: $\pi_{SE}^* = \arg \max_{\pi} \sum_a P^{\pi} E_{S \sim \pi} [r(S,a) + \alpha \mathcal{H}(\pi(s))]$

value function: $V_{SE}^{\pi}(s) = \sum_a P^{\pi} E_{S \sim \pi} [r(S,a) + \alpha \mathcal{H}(\pi(s))]$

Action-value function: $Q_{SE}^{\pi}(s,a) = r(s,a) + \sum_{S'} P(S'|s,a) E_{S' \sim \pi} [r(S',a) + \alpha \mathcal{H}(\pi(S'))]$

Theorem

$$Q_{SE}^{\pi}(s,a) = r(s,a) + \gamma E_s [V_{SE}^{\pi}(s')]$$

$$V_{SE}^{\pi}(s) = \sum_a \pi(a) Q_{SE}^{\pi}(s,a) \quad \text{Easy to see}$$

Theorem: Energy Form

Given a policy π with value function $Q_{SE}^{\pi}, V_{SE}^{\pi}$
the greedy policy wrt π is $\tilde{\pi}(s) = \arg \max_{\pi} \sum_a \tilde{\pi}(a) Q_{SE}^{\pi}(s,a) + \alpha \mathcal{H}(\tilde{\pi}(s))$

$$Q_{SE}^{\pi}(s, \tilde{\pi}(s)) = \alpha \log \sum_a \exp(Q_{SE}^{\pi}(s,a) / \alpha)$$

$$\tilde{\pi}(a) = \frac{\exp(Q_{SE}^{\pi}(s,a) / \alpha)}{\sum_b \exp(Q_{SE}^{\pi}(s,b) / \alpha)} = \exp\left[\frac{1}{\alpha} [Q_{SE}^{\pi}(s,a) - Q_{SE}^{\pi}(s, \tilde{\pi}(s))]\right]$$

$$\tilde{\pi}(s) = \arg \max_{\pi} \sum_a \pi(a) Q_{SE}^{\pi}(s,a) + \mathcal{H}(\pi(s))$$

$$= \arg \max_{\pi} \sum_a \pi(a) [Q_{SE}^{\pi}(s,a) - \alpha \log \pi(a)]$$

$$\min_{\pi} \left\{ - \sum_i \pi_i \left(a_i - \alpha \log \pi_i \right) \right\} \Rightarrow \text{convex} \quad \left(\begin{array}{l} \text{optimal value is negated} \\ \text{when flip max to min} \end{array} \right)$$

$$\mathcal{L}(x, v) = -\sum_i x_i (a_i - \alpha \log x_i) + v \left(\sum_i x_i - 1 \right)$$

$$\frac{\partial \mathcal{L}}{\partial x_i} = -a_i + \alpha (1 + \log x_i) + v \stackrel{!}{=} 0$$

$$x_i^*(v) = \exp \left(\frac{a_i - v}{\alpha} - 1 \right)$$

$$\begin{aligned} g(v) &= -\sum_i \exp \left(\frac{a_i - v}{\alpha} - 1 \right) (v + \alpha) + v \left(\sum_i \exp \left(\frac{a_i - v}{\alpha} - 1 \right) - 1 \right) \\ &= -\alpha \sum_i \exp \left(\frac{a_i - v}{\alpha} - 1 \right) - v \end{aligned}$$

$$\frac{dg}{dv} = \sum_i \exp \left(\frac{a_i - v}{\alpha} - 1 \right) - 1 = \exp \left(\frac{-v}{\alpha} \right) \sum_i \exp \left(\frac{a_i}{\alpha} - 1 \right) - 1 \stackrel{!}{=} 0$$

$$v^* = \alpha \log \sum_i \exp \left(\frac{a_i}{\alpha} - 1 \right)$$

$$x_i^* = \frac{\exp \left(\frac{a_i}{\alpha} - 1 \right)}{\exp \left(\frac{-v^*}{\alpha} \right)} = \frac{\exp \left(\frac{a_i}{\alpha} - 1 \right)}{\sum_j \exp \left(\frac{a_j}{\alpha} - 1 \right)} = \frac{\exp \left(\frac{a_i}{\alpha} \right)}{\sum_j \exp \left(\frac{a_j}{\alpha} \right)}$$

$$p^* = -\sum_i \frac{\exp(a_i/k)}{\sum_j \exp(a_j/k)} \left(a_i - \alpha \log \frac{\exp(a_i/k)}{\sum_j \exp(a_j/k)} \right)$$

$$= -\frac{1}{\alpha} \sum_i \exp(a_i/k) (\alpha \log z)$$

$$= -\alpha \log z$$

$$= -\alpha \log \sum_j \exp(a_j/k)$$

$$p^* = -\alpha \log \sum_j \exp(a_j/k)$$

$$\longrightarrow x_i^* = \exp \left(\frac{a_i}{\alpha} + \frac{p^*}{\alpha} \right)$$

Theorem: Monotonic Improvement

given a policy π with value function $Q_{\pi}^{\pi_0}$

let one-step greedy policy $\tilde{\pi}(a|s) = \frac{\exp(Q_{\pi}^{\tilde{\pi}}(s,a)/k)}{\sum_a \exp(Q_{\pi}^{\tilde{\pi}}(s,a)/k)}$

then $Q_{\pi}^{\tilde{\pi}}(s,a) \geq Q_{\pi}^{\pi}(s,a) \forall (s,a)$

Thus the optimal policy must have form $\pi^*(a|s) = \frac{\exp(Q_{\pi^*}^{\pi^*}(s,a)/k)}{\sum_a \exp(Q_{\pi^*}^{\pi^*}(s,a)/k)}$

$$\begin{aligned} Q_{\pi^*}^{\pi^*}(s, a) &= E_{s_1} [r_0 + \gamma (\alpha \mathcal{H}(\pi^*(s_1)) + E_{a_1, \pi^*} [Q_{\pi^*}^{\pi^*}(s_1, a_1)])] \\ &\leq E_{s_1} [r_0 + \gamma (\alpha \mathcal{H}(\tilde{\pi}(s_1)) + E_{a_1, \tilde{\pi}} [Q_{\pi^*}^{\pi^*}(s_1, a_1)])] \\ &\leq \dots \\ &\leq Q_{\pi^*}^{\tilde{\pi}}(s, a) \end{aligned}$$

Theorem: Bellman Optimality

By the monotonic improvement theorem $\pi^*(a|s) = \frac{\exp(Q_{\pi^*}^*(s,a))}{\sum_a \exp(Q_{\pi^*}^*(s,a))}$

then $\pi^* = \arg\max_{\pi} \sum_a \pi(a|s) Q_{\pi}^*(s,a) + \alpha \mathcal{H}(\pi(s))$

$$\begin{aligned} V_{\pi^*}^*(s) &= \sum_a \pi^*(a|s) Q_{\pi^*}^*(s,a) + \alpha \mathcal{H}(\pi^*(s)) \\ &= \max_{\pi} \sum_a \pi(a|s) Q_{\pi}^*(s,a) + \alpha \mathcal{H}(\pi(s)) \\ &= \alpha \log \sum_a \exp(Q_{\pi^*}^*(s,a) / \alpha) \end{aligned}$$

$$Q_{\pi^*}^*(s,a) = r(s,a) + \gamma E_{s'} [V_{\pi^*}^*(s')]$$

$$= r(s,a) + \gamma E_{s'} [\alpha \log \sum_a \exp(Q_{\pi^*}^*(s',a) / \alpha)]$$

at each state, the optimal policy π^* is greedy wrt $Q_{\pi^*}^*$

Theorem: Bellman Backup and uniqueness of π^*

define the soft value iteration as

$$(TQ)(s,a) = r(s,a) + \gamma E_{s'} [\alpha \log \int_a \exp(Q_{\pi^*}^*(s',a') / \alpha) da']$$

Then T is a contraction

thus there's a unique $Q_{\pi^*}^*$ satisfying

$$Q_{\pi^*}^*(s,a) = r(s,a) + \gamma E_{s'} [\alpha \log \int_a \exp(Q_{\pi^*}^*(s',a') / \alpha) da']$$

and the optimal policy is $\pi^*(a|s) = \frac{\exp(Q_{\pi^*}^*(s,a) / \alpha)}{\int_a \exp(Q_{\pi^*}^*(s,a') / \alpha) da'}$

$$\text{let } \|Q_1 - Q_2\|_0 = \max_{s,a} |Q_1(s,a) - Q_2(s,a)| = \epsilon$$

$$\begin{aligned} \alpha \log \int_a \exp \frac{Q_1(s,a)}{\alpha} da &\leq \alpha \log \int_a \exp \frac{Q_2(s,a) + \epsilon}{\alpha} da \\ &= \alpha \log \int_a \exp \frac{\epsilon}{\alpha} \exp \frac{Q_2(s,a)}{\alpha} da \\ &= \epsilon + \alpha \log \int_a \exp \frac{Q_2(s,a)}{\alpha} da \quad \text{HS} \end{aligned}$$

$$\alpha \log \int_a \exp \frac{Q_1(s,a)}{\alpha} da \geq -\epsilon + \alpha \log \int_a \exp \frac{Q_2(s,a)}{\alpha} da \quad \text{HS}$$

$$-\epsilon \leq \alpha \log \int_A \exp \frac{Q_1(s,a)}{\alpha} da - \alpha \log \int_A \exp \frac{Q_2(s,a)}{\alpha} da \leq \epsilon \quad \forall s$$

$$\begin{aligned} \|TQ_1 - TQ_2\|_\infty &= \max_{s,a} \left| E_{s'} \left[\alpha \log \int_A \exp \frac{Q_1(s,a)}{\alpha} da \right] - E_{s'} \left[\alpha \log \int_A \exp \frac{Q_2(s,a)}{\alpha} da \right] \right| \\ &= \max_{s,a} \left| E_{s'} \left[\alpha \log \int_A \exp \frac{Q_1(s,a)}{\alpha} da - \alpha \log \int_A \exp \frac{Q_2(s,a)}{\alpha} da \right] \right| \\ &\leq \epsilon \\ &= \epsilon \|Q_1 - Q_2\|_\infty \end{aligned}$$

• Soft Q-Iteration

$$\begin{aligned} Q_{\text{soft}}(s,a) &\leftarrow r(s,a) + \gamma \cdot \alpha \log \int_A \exp \frac{Q(s,a')}{\alpha} da' \quad \xrightarrow{\alpha \log \int_A \exp \frac{Q(s,a')}{\alpha} da' \text{ for discrete action space}} \\ &= r(s,a) + \gamma \alpha \log E_{g(a')} \left[\frac{Q(s,a')}{g(a')} \right] \end{aligned}$$

initialize Q^θ and target $Q^\bar{\theta} = Q^\theta$

initialize replay buffer D

while true

sample $(s,a,r,s') \sim D$ g can be any distribution, for example uniform

sample $\alpha_i \sim g(a')$

compute $t = r + \gamma \log \left[\frac{1}{m} \sum_{i=1}^m \frac{\exp(Q^\bar{\theta}(s',a_i)/\alpha_i)}{g(a_i)} \right]$ but scales poorly to high dimension
one choice is $g(a|s) = \frac{\exp(Q(s,a)/\alpha)}{\int_a \exp(Q(s,a)/\alpha) da}$

$\min_\theta \|Q^\theta(s,a) - t\|_2^2$

for each k iteration: $\bar{\theta} \leftarrow \theta$

but this will require SGD
for continuous action space

Soft Actor-Critic (Standard)

Compute Q_θ π_ϕ V_ψ

$$\text{Policy } \pi_\phi(a|s) = \frac{\exp(Q_\phi(s,a)/\alpha)}{\int_a \exp(Q_\phi(s,a)/\alpha) da}$$

$$\min_{\phi} D_{KL}(\pi_\phi(\cdot|s) \parallel \frac{\exp(Q_\phi(s,a)/\alpha)}{\int_a \exp(Q_\phi(s,a)/\alpha) da})$$

$$\min_{\phi} E_{a \sim \pi_\phi(s)} [\log \pi_\phi(a|s) - Q_\phi(s,a)/\alpha]$$

$$\begin{aligned} &= \min_{\phi} \int_a \pi_\phi(a|s) \cdot \log \frac{\pi_\phi(a|s)}{\exp(Q_\phi(s,a)/\alpha) / Z_\phi(s)} da \\ &= \min_{\phi} E_{a \sim \pi_\phi(s)} [\log \pi_\phi(a|s) - Q_\phi(s,a)/\alpha + \log Z_\phi(s)] \\ &= \min_{\phi} E_{a \sim \pi_\phi(s)} [\log \pi_\phi(a|s) - Q_\phi(s,a)/\alpha] \end{aligned}$$

let $\pi_\phi(s) = U_\phi(s) + \phi_\phi(s) \cdot \epsilon$ $\epsilon \sim \mathcal{N}(0, I)$

$$\min_{\phi} E_{\epsilon \sim \mathcal{N}} [\alpha \log \pi_\phi(U_\phi(s) + \phi_\phi(s) \cdot \epsilon) - Q_\phi(s, U_\phi(s) + \phi_\phi(s) \cdot \epsilon)]$$

Initialize $V_\psi, V_\phi = V_\psi, Q_\theta, Q_\phi, \pi_\phi$

while true {

Sample $(s, a, r, s') \sim \mathcal{D}$

learn Q to approximate $Q_{soft}^*(s, a) = r(s, a) + \gamma E_{s'} [V_{soft}^*(s')]$

compute target $y(s) = r(s, a) + \gamma V_\psi(s')$

$$\min_{\theta} \|Q_\theta(s, a) - y(s)\|_2^2 \quad \min_{\phi} \|Q_\phi(s, a) - y(s)\|_2^2$$

Sample $\epsilon \sim \mathcal{N}(0, I)$ $\bar{a} = U_\phi(s) + \phi_\phi(s) \cdot \epsilon$

learn V to approximate $V_{soft}^*(s) = \alpha \mathcal{H}(\pi^*(s)) + \int_a \pi(a|s) \cdot Q_{soft}^*(s, a) da$

$$= E_{a \sim \pi_\phi} [Q_{soft}^*(s, a) - \alpha \log \pi^*(a|s)]$$

compute target $y(s) = \min(Q_\theta(s, \bar{a}), Q_\phi(s, \bar{a})) - \alpha \log \pi^*(a|s)$

$$\min_{\psi} \|V_\psi(s) - y(s)\|_2^2$$

learn π to approximate $\pi^*(a|s) = \frac{\exp(Q^*(s, a)/\alpha)}{\int_a \exp(Q^*(s, a)/\alpha) da}$

$$\min_{\phi} \alpha \log \pi_\phi(a|s) - \min(Q_\theta(s, \bar{a}), Q_\phi(s, \bar{a}))$$

$$\bar{\Psi} = \rho \bar{\Psi} + (1-\rho) \Psi$$

}

Soft Actor-Critic (Modern)

Initialize $Q_{\bar{\theta}_1} = Q_{\theta_1}$ $Q_{\bar{\theta}_2} = Q_{\theta_2}$ π_0

while true {

sample $(S, a, r, s') \sim D$

$$\begin{aligned} \text{learn } Q \text{ to approximate } Q_{\text{soft}}^*(s, a) &= r(s, a) + \gamma E_{s'} [V_{\text{soft}}^*(s')] \\ &= r(s, a) + \gamma E_{s'} \left[\int_{a'} \pi^*(a'|s) Q_{\text{soft}}^*(s', a') + \alpha \log \pi^*(a'|s) \right] \\ &= r(s, a) + \gamma E_{s'} E_{a' \sim \pi(s)} [Q_{\text{soft}}^*(s', a') - \alpha \log \pi^*(a'|s)] \end{aligned}$$

sample $a' \sim \pi_\phi(a'|s) \sim U_\phi(s') + b_\phi(s') + b$

compute target $y(s) = r(s, a) + \gamma [\min(\bar{Q}_{\bar{\theta}_1}(s', a'), \bar{Q}_{\bar{\theta}_2}(s', a')) - \alpha \log \pi^*(a'|s)]$

$$\min_{\theta_1} \|Q_{\theta_1}(s, a) - y(s)\|_2^2 \quad \min_{\theta_2} \|Q_{\theta_2}(s, a) - y(s)\|_2^2$$

$$\text{learn } \pi \text{ to approximate } \pi^*(a|s) = \frac{\exp(Q^*(s, a)/\alpha)}{\int_a \exp(Q^*(s, a)/\alpha) da}$$

sample $a \sim \pi_\phi(a|s)$

$$\min_{\phi} \int \log \pi_\phi(a|s) - \min(\bar{Q}_{\bar{\theta}_1}(s, \bar{a}), \bar{Q}_{\bar{\theta}_2}(s, \bar{a}))$$

$$\bar{\theta}_1 = \rho \bar{\theta}_1 + (1-\rho) \theta_1$$

$$\bar{\theta}_2 = \rho \bar{\theta}_2 + (1-\rho) \theta_2$$

}

• Automating Entropy Adjustment

An Intuition (Based on dual gradient method explanation in EE364B)
the dual gradient is basically adaptive penalty for violation of constraint.)

Constraint: $E_{a \sim \pi_A} [\log \pi(a|s)] \leq -\bar{\mathcal{H}} \quad : \alpha$

update $\alpha \leftarrow \underbrace{E_{a \sim \pi_A} [\log \pi(a|s)] + \bar{\mathcal{H}}}_{\text{violation}} \rightarrow \bar{\mathcal{H}} = -\dim A$

if this is positive, the constraint is violated
then penalty is increased