

VALUE ITERATION, POLICY ITERATION: PRELIMINARY RESULTS

VI-P1. 1

Recall

$$V^*(x) = \min_{\mu} \left[\sum_{t=0}^{+\infty} \gamma^t \mathbb{E} [g(x_t, \mu(x_t), w_t)] \right] \quad x_0 = x$$

optimal policy is stationary

$$V^\mu(x) = \sum_{t=0}^{+\infty} \gamma^t \mathbb{E} [g(x_t, \mu(x_t), w_t)]$$

$$\min_{\mu} \mathbb{E} [g(x, \mu, w) + \gamma V^*(f(x, \mu, w))]$$

we know that (Bellman's equation)

$$V^*(x) = \min_{\mu} \left[\mathbb{E} [g(x, \mu, w_0)] + \gamma \sum_{i=1}^M P(i|x, \mu) V^*(i) \right] \quad \forall x \in X$$

$$\Rightarrow T^*[V] : \mathbb{R}^M \rightarrow \mathbb{R}^M \quad \min_{\mu} \mathbb{E} [g(x, \mu, w) + \gamma V(f(x, \mu, w))]$$

$$T^*[V]_x = \min_{\mu} \left[\mathbb{E} [g(x, \mu, w_0)] + \gamma \sum_{i=1}^M P(i|x, \mu) V(i) \right] \quad \forall x \in X$$

$$V^* = \text{solution to } V = T^*[V] \quad (\text{fixed point})$$

similarly, $\forall x \in X$

$$V^\mu(x) = \sum_{t=0}^{+\infty} \gamma^t \mathbb{E} [g(x_t, \mu(x_t), w_t)] = \mathbb{E} [g(x, \mu(x), w_0)] + \gamma \sum_{t=1}^{+\infty} \gamma^{t-1} \mathbb{E} [g(x_t, \mu(x_t), w_t)]$$

$\tau = t-1$

$$= \mathbb{E} [g(x, \mu(x), w_0)] + \gamma V^\mu(x_1) = \mathbb{E} [g(x, \mu(x), w)] + \gamma \sum_{i=1}^M P(i|x, \mu(x)) V^\mu(i)$$

$$\Rightarrow T^\mu[V] : \mathbb{R}^M \rightarrow \mathbb{R}^M \quad T^\mu[V]_x = \mathbb{E} [g(x, \mu(x), w_0)] + \gamma \sum_{i=1}^M P(i|x, \mu(x)) V(i) \quad x \in X$$

$$V^\mu = \text{solution to } V = T^\mu[V] \quad (\text{fixed point}) \quad \mathbb{E} [g(x, \mu(x), w)] + \gamma V(f(x, \mu(x), w))$$

result $V = T^*[V]$ and $V = T^\mu[V]$... solution exist and is

unique $\Leftarrow$ Banach-Caccioppoli theorem (the tool)

(S, d) = ^{complete} metric space $F: S \rightarrow S$ contraction mapping

$$\text{i.e. } \exists \rho < 1 : d(F(s_1), F(s_2)) \leq \rho \cdot d(s_1, s_2) \quad \forall s_1, s_2 \in S$$

$s = F(s)$ always admits one and only one solution

clearly this implies continuity

Proof: uniqueness

VI-PI.2

By contradiction, suppose that $s_1 = F(s_1)$ and $s_2 = F(s_2)$. Then

$$d(s_1, s_2) = d(F(s_1), F(s_2)) \leq p d(s_1, s_2) < d(s_1, s_2) \quad \dots \text{contradiction!}$$

$p < 1$

existence

let s_0 be an arbitrary point in S and define

$$s_{n+1} = F(s_n) = F^{n+1}(s_0) \quad n = 0, 1, 2, \dots$$

$$\text{Then } d(s_{n+1}, s_n) = d(F(s_n), F(s_{n-1})) \leq p d(s_n, s_{n-1}) \leq \dots \leq p^n d(s_1, s_0)$$

For any $m > n$
triangle ineq.

$$d(s_m, s_n) \leq d(s_m, s_{m-1}) + d(s_{m-1}, s_{m-2}) + \dots + d(s_{n+1}, s_n) \leq \dots$$

$$\dots \leq \sum_{k=n}^{m-1} d(s_{k+1}, s_k) \leq \sum_{k=n}^{m-1} p^k d(s_1, s_0) \leq \sum_{k=n}^{\infty} p^k d(s_1, s_0) = \frac{p^n}{1-p} d(s_1, s_0)$$

so for any $m > n$, $d(s_m, s_n) \xrightarrow{m, n \rightarrow \infty} 0$, i.e. $\{s_m\}$ is Cauchy convergent

$$\text{completeness} \Rightarrow \exists \bar{s} : s_m \xrightarrow{m \rightarrow \infty} \bar{s}$$

$$\text{However, } \bar{s} = \lim_{n \rightarrow \infty} s_{n+1} = \lim_{n \rightarrow \infty} F(s_n) = F(\bar{s}) \quad \text{i.e. } \bar{s} \text{ is a fixed point.}$$

QED

$s_0 = \text{arbitrary}$

$s_{n+1} = F(s_n) \Rightarrow$ fixed-point iteration
(subsequent approximations)

$\rightarrow$ algorithm to find fixed point
 $\bar{s} = F(\bar{s})$

$$d(s_{n+1}, \bar{s}) = d(F(s_n), F(\bar{s})) \leq p d(s_n, \bar{s}) \leq \dots \leq p^{n+1} d(s_0, \bar{s}) \quad (\text{exponential convergence})$$

Fundamental property

Both T^* and T^M are contractions from $\mathbb{R}^M$ to $\mathbb{R}^M$ with respect to the max-norm distance

T^* : let V_1 and V_2 two value functions (i.e. two vectors in $\mathbb{R}^M$)

$$\max_{x \in X} \left| T^*[V_1] - T^*[V_2] \right| = \max_{x \in X} \left| \min_u \left[E[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) V_1(i) \right] + \right. \\ \left. - \min_u \left[E[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) V_2(i) \right] \right| \leq$$

VI-PI-5
DP-3
≤ assume $\min_u [\dots V_2(i)] \leq \min_u [\dots V_1(i)]$ and attained for u_2^*

(if $\min_u [\dots V_1(i)] \leq \min_u [\dots V_2(i)]$ evaluate everything in u_1^* where $\min_u [\dots V_1(i)]$ is attained)
 $\rightarrow$ distance increases $\leftarrow \min_u$

$$\leq \max_{x \in X} \left| \mathbb{E}[g(x, u_2^*, w)] + \gamma \sum_{i=1}^m p(i|x, u_2^*) V_1(i) - \mathbb{E}[g(x, u_2^*, w)] - \gamma \sum_{i=1}^m p(i|x, u_2^*) V_2(i) \right|$$

$$= \max_{x \in X} \left| \gamma \sum_{i=1}^m p(i|x, u_2^*) (V_1(i) - V_2(i)) \right| \leq \max_{x \in X} \gamma \sum_{i=1}^m p(i|x, u_2^*) |V_1(i) - V_2(i)|$$

$$\leq \max_{i \in X} |V_1(i) - V_2(i)| \cdot \max_{x \in X} \gamma \sum_{i=1}^m p(i|x, u_2^*) =$$

$$= \gamma \cdot \max_{x \in X} |V_1(x) - V_2(x)|$$

T^μ : same, simpler V_1 and $V_2 \in \mathbb{R}^m$ value functions (ie)

$$\max_{x \in X} |T^\mu[V_1] - T^\mu[V_2]| =$$

$$= \max_{x \in X} \left| \mathbb{E}[g(x, \mu(x), w)] + \gamma \sum_{i=1}^m p(i|x, \mu(x)) V_1(i) - \mathbb{E}[g(x, \mu(x), w)] - \gamma \sum_{i=1}^m p(i|x, \mu(x)) V_2(i) \right|$$

$$= \max_{x \in X} \left| \gamma \sum_{i=1}^m p(i|x, \mu(x)) (V_1(i) - V_2(i)) \right| \leq$$

$$\leq \gamma \max_{x \in X} \sum_{i=1}^m p(i|x, \mu(x)) |V_1(i) - V_2(i)| \leq \gamma \max_{i \in X} |V_1(i) - V_2(i)| \cdot \max_{x \in X} \sum_{i=1}^m p(i|x, \mu(x)) =$$

$$= \gamma \max_{x \in X} |V_1(x) - V_2(x)|$$

Thus, $V = T^*[V]$ and $V = T^\mu[V]$ admits one and only one solution $\leadsto V^*(x)$ and $V^\mu(x)$

optimal policy $\iff$ optimal value function $V^*(x)$ } key problem: compute $V^*(x), V^\mu(x)$
 policy iteration $\iff$ value function for μ $V^\mu(x)$

VALUE ITERATION AND POLICY ITERATION ALGORITHMS

$$V^*(x) = \min_u \left[\mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^m p(i|x, u) V^*(i) \right] \quad x=1, 2, \dots, m \quad \text{VI-PI.4}$$

piecewise linear in V^* $\Rightarrow$ system of non-linear eqs

$$V^\mu(x) = \mathbb{E}[g(x, \mu(x), w)] + \gamma \sum_{i=1}^m p(i|x, \mu(x)) V^\mu(i)$$

$\Downarrow$ linear in $V^\mu \Rightarrow$ system of linear eqs

$(I - \gamma \cdot P) V^\mu = [\mathbb{E}[g(x, \mu(x), w)]]_{x=1}^m$ better but still tough if m is large

tough (impossible, curse of dimensionality)

$\Rightarrow$ use fixed-point iterations! $\Rightarrow$ VALUE ITERATION algorithm (VI)

$V^0(x) \quad x \in X$ arbitrarily chosen (need not be a value function associated to some policy)

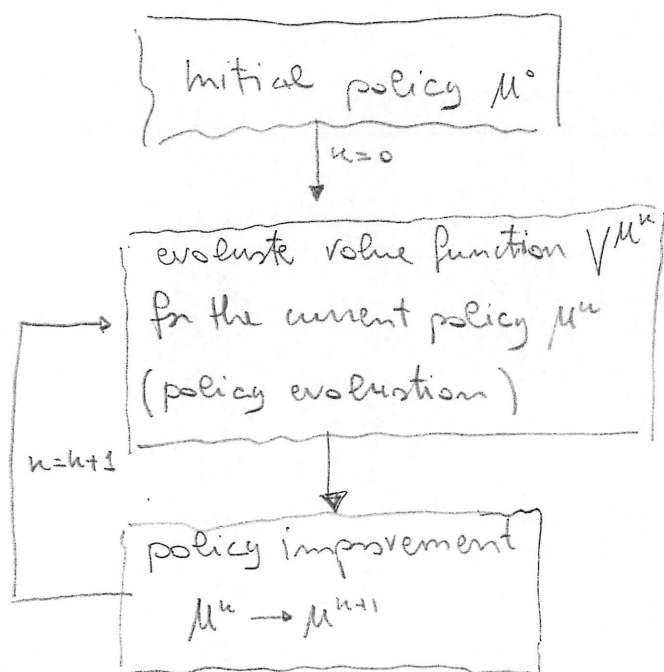
$$V^{k+1}(x) = T^*[V^k(x)] \text{ converges to } V^*(x)$$

$$V^{k+1}(x) = T^\mu[V^k(x)] \text{ converges to } V^\mu(x)$$

exponentially fast

OBS: $V^k(x) = [T^*]^k [V^0(x)] \quad \text{or} \quad V^k(x) = [T^\mu]^k [V^0(x)]$

Alternative to VI to compute $V^* \Rightarrow$ POLICY ITERATION (PI)



• sequence of policies, improving at each iteration ($V^{\mu^{n+1}} \leq V^{\mu^n} \forall x \in X$)

• convergence in a finite no. of iterations (can be very large)

• faster convergence than VI for large state space

big improvement in the first few iterations

property: MONOTONICITY

VI-PI.4.bis

Suppose $V_1(x) \leq V_2(x) \quad \forall x \in X$. Then, for every μ we have

$$T^\mu[V_1]_x \leq T^\mu[V_2]_x \quad \forall x \in X \quad \text{and} \quad T^*[V_1]_x \leq T^*[V_2]_x \quad \forall x \in X$$

Proof:

$$\begin{aligned} T^\mu[V_1]_x &= \mathbb{E}[g(x, \mu(x), w)] + \gamma \sum_{i=1}^M p(i|x, \mu(x)) V_1(i) \\ &\leq \mathbb{E}[g(x, \mu(x), w)] + \gamma \sum_{i=1}^M p(i|x, \mu(x)) V_2(i) = T^\mu[V_2]_x \end{aligned}$$

$$\mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M \overset{\geq 0}{p(i|x, u)} V_1(i) \leq \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M \overset{\geq 0}{p(i|x, u)} V_2(i)$$

$\forall x \in X \quad \forall u \in \mathcal{U}$. Hence,

$$\begin{aligned} T^*[V_1]_x &= \min_{u \in \mathcal{U}} \left[\mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M p(i|x, u) V_1(i) \right] \\ &\leq \min_{u \in \mathcal{U}} \left[\mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M p(i|x, u) V_2(i) \right] = T^*[V_2]_x \end{aligned}$$

Start with an initial policy $\mu^0: \mathcal{X} \rightarrow \mathcal{U}$ arbitrary

VI-Pl.5

Policy evolution: solve $V^{\mu^k}(x) = \mathbb{E}[g(x, \mu^k(x), w)] + \gamma \sum_{i=1}^M p(i|x, \mu^k(x)) \cdot V^{\mu^k}(i)$
 $\forall x \in \mathcal{X}$

$\hookrightarrow$ either solve system of linear eqs. or use VI for V^{μ^k} exact full convergence $\leadsto$ solve $V = T^{\mu^k}[V]$

Policy improvement: $\mu^{k+1}(x) = \underset{u \in \mathcal{U}}{\operatorname{argmin}} \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M p(i|x, u) V^{\mu^k}(i)$

Properties:

$$1. V^{\mu^{k+1}}(x) \leq V^{\mu^k}(x) \quad \forall x \in \mathcal{X} \quad / \quad V^0 = T^{\mu^{k+1}}[V^{\mu^k}]$$

$$\begin{aligned} \text{Define } \forall x \in \mathcal{X}, \quad V^0(x) &= \mathbb{E}[g(x, \mu^{k+1}(x), w)] + \gamma \sum_{i=1}^M p(i|x, \mu^{k+1}(x)) \cdot V^{\mu^k}(i) \\ &= \min_{u \in \mathcal{U}} \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M p(i|x, u) V^{\mu^k}(i) \\ &\leq \mathbb{E}[g(x, \mu^k(x), w)] + \gamma \sum_{i=1}^M p(i|x, \mu^k(x)) V^{\mu^k}(i) = V^{\mu^k}(x) \end{aligned}$$

$$\text{so } V^0(x) \leq V^{\mu^k}(x) \quad \forall x \in \mathcal{X}$$

Then

$$V^{h+1} = T^{\mu^{k+1}}[V^h] \quad h=0, 1, \dots \quad \xRightarrow{\text{VI}} \quad V^h \xrightarrow{h \rightarrow \infty} V^{\mu^{k+1}}$$

$$V^1 = T^{\mu^{k+1}}[V^0] \leq (\text{monotonicity, } V^0 \leq V^{\mu^k}) \leq T^{\mu^{k+1}}[V^{\mu^k}] = V^0$$

$$V^2 = T^{\mu^{k+1}}[V^1] \leq (\text{monotonicity, } V^1 \leq V^0) T^{\mu^{k+1}}[V^0] = V^1 \leq V^0$$

$\vdots$

$$V^h \leq V^0 \leq V^{\mu^k} \Rightarrow V^{\mu^{k+1}} = \lim_{h \rightarrow \infty} V^h \leq V^{\mu^k}$$

2. no. of ^{stationary} policies is finite ($\leq m^m$) $\Rightarrow V^{\mu^{k+1}}(x) < V^{\mu^k}(x)$ for some x

occurs for a finite no. of iterations, then for some k

$$V^{\mu^{k+1}}(x) = V^{\mu^k}(x) \quad \forall x \in \mathcal{X} \quad \rightarrow \quad V^{\mu^{k+1}} = V^h \quad \forall h \text{ in previous iterations}$$

$$\Rightarrow V^{\mu^k}(x) = V^{\mu^{k+1}}(x) = V^0(x) = \min_u \left[\mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M p(i|x, u) V^{\mu^k}(i) \right]$$

i.e. $V^{\mu^n}(x)$ satisfies the Bellman eq. for optimal value $V = P1.6$
 function $\Rightarrow \mu^n = \mu^*$ optimal policy, $V^{\mu^n} = V^*$ optimal
 value function

PI converges in finite time to the optimal policy!

Q-FACTORS (action-value functions)

$$Q^*: X \times U \rightarrow \mathbb{R} \quad Q^\mu: X \times U \rightarrow \mathbb{R}$$

$$Q^*(x, u) = \mathbb{E}[g(x, u, w_0)] + \gamma \sum_{i=1}^M P(i|x, u) V^*(i) \quad x \in X \quad u \in U$$

$$Q^\mu(x, u) = \mathbb{E}[g(x, u, w_0)] + \gamma \sum_{i=1}^M P(i|x, u) V^\mu(i) \quad x \in X \quad u \in U$$

interpretation: total discounted cost achieved when at $t=0$ (initial time)

$x_0 = x$ and $u_t = u$ (action u is used) and then.

$*$: optimal policy μ^* is applied for $t \geq 1$

μ : policy μ is applied for $t \geq 1$

Clearly, thanks to Bellman equations

$$\bullet \quad V^*(x) = \min_{u \in U} Q^*(x, u) \quad \mu^*(x) = \underset{u \in U}{\operatorname{argmin}} Q^*(x, u)$$

$$\bullet \quad V^\mu(x) = Q^\mu(x, \mu(x)) \quad Q^* \text{ summarizes all the info we need to determine the optimal policy}$$

Bellman equations for Q-factors (from definitions and relations above)

$$Q^*(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) \cdot \min_{v \in U} Q^*(i, v)$$

$$\text{solution to } Q = T^*[Q] \text{ where } T^*[Q] = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) \min_v Q(i, v)$$

$$Q^\mu(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) Q^\mu(i, \mu(i))$$

$$\text{solution to } Q = T^\mu[Q] \text{ where } T^\mu[Q] = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) Q(i, \mu(i))$$

property: T^* and T^M (for Q-factors) are contractions with respect to the $\max_{x,u}$ distance

$Q_1(x,u), Q_2(x,u)$ given.

$$\begin{aligned} \max_{x,u} |T^*[Q_1] - T^*[Q_2]| &= \max_{x,u} \left| \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) \min_v Q_1(i,v) - \mathbb{E}[g(x,u,w)] + \right. \\ &\quad \left. - \gamma \sum_{i=1}^M P(i|x,u) \min_v Q_2(i,v) \right| \leq \max_{x,u} \left| \gamma \sum_{i=1}^M P(i|x,u) \left| \min_v Q_1(i,v) - \min_v Q_2(i,v) \right| \right| \\ &\leq \max_{x,u} \left| \gamma \sum_{i=1}^M P(i|x,u) |Q_1(i, v(i)) - Q_2(i, v(i))| \right| \quad v(i) = \begin{cases} \arg \min_v Q_1(i,v) & \text{if } \min_u Q_1(i,u) \leq \min_u Q_2(i,u) \\ \arg \min_v Q_2(i,v) & \text{else} \end{cases} \\ &\leq \max_{i,v} |Q_1(i,v) - Q_2(i,v)| \cdot \max_{x,u} \gamma \sum_{i=1}^M P(i|x,u) \\ &= \gamma \cdot \max_{x,u} |Q_1(x,u) - Q_2(x,u)| \end{aligned}$$

$$\begin{aligned} \max_{x,u} |T^M[Q_1] - T^M[Q_2]| &= \max_{x,u} \left| \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) Q_1(i, \mu(i)) - \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) Q_2(i, \mu(i)) \right| \\ &\leq \max_{x,u} \gamma \cdot \sum_{i=1}^M P(i|x,u) |Q_1(i, \mu(i)) - Q_2(i, \mu(i))| \stackrel{\text{as before}}{\leq} \dots \leq \gamma \max_{x,u} |Q_1(x,u) - Q_2(x,u)| \end{aligned}$$

property: T^* and T^M (for Q-factors) are monotonic

$Q_1(x,u) \leq Q_2(x,u) \quad \forall x \in X, u \in U$ Then, $\forall x \in X, u \in U$

$$\begin{aligned} T^*[Q_1](x,u) &= \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) \min_v Q_1(i,v) \\ &\leq \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) \min_v Q_2(i,v) = T^*[Q_2](x,u) \end{aligned}$$

$$\begin{aligned} T^M[Q_1](x,u) &= \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) Q_1(i, \mu(i)) \\ &\leq \mathbb{E}[g(x,u,w)] + \gamma \sum_{i=1}^M P(i|x,u) Q_2(i, \mu(i)) = T^M[Q_2](x,u) \end{aligned}$$

FACT: since T^* and T^M are contractions, the Bellman eqs. for Q-factors univocally determine $Q^*(x,u)$ and $Q^M(x,u)$ always (can be proved otherwise) and shows the validity of the following Value Iteration method for Q-factors.

in a model-based environment $\rightarrow V^*$ and Q^* are equivalent
 in a model-free (RL):

VI-14:

$$\mu^*(u) = \underset{u}{\operatorname{argmin}} Q^*(x, u) = \underset{u}{\operatorname{argmin}} \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) V^*(i)$$

$\uparrow$ that's all we need to compute the optimal policy $\uparrow$ model needed besides V^*

Moreover, $\mathbb{E}[g(x, u, w)] =$

$$Q^*(x, u) = T^*[Q^*](x, u) = \mathbb{E}[g(x, u, w)] + \gamma \min_v Q^*(x, u, w, v)$$

$\uparrow$ expected value outside $\uparrow$ min inside $\Rightarrow$ suitable for incremental computation

$$V^*(x) = T^*[V^*](x) = \min_u [\mathbb{E}[g(x, u, w)] + \gamma V^*(P(x, u, w))]$$

$\uparrow$ outside $\uparrow$ inside $\rightarrow$ not suitable for incremental computation

VALUE ITERATION for Q-factors

$$Q^0(x, u) : \mathcal{X} \times \mathcal{U} \rightarrow \mathbb{R} \text{ arbitrary}$$

$$Q^{k+1}(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) \cdot \min_v Q^k(i, v) \xrightarrow{k \rightarrow \infty} Q^*(x, u)$$

$$= T^*[Q^k](x, u)$$

$$Q^{k+1}(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) Q^k(i, \mu(i)) \xrightarrow{k \rightarrow \infty} Q^\mu(x, u)$$

$$= T^\mu[Q^k](x, u)$$

either solve sys. of linear equations or use VI

POLICY ITERATION for Q-factors

$$\mu^0(x) : \mathcal{X} \rightarrow \mathcal{U} \text{ initial policy (arbitrary)}$$

policy evaluation: solve $Q^{\mu^k}(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) Q^{\mu^k}(i, \mu^k(i))$

$$Q^{\mu^k}(x, u) = T^{\mu^k}[Q^{\mu^k}](x, u) \quad \forall x \in \mathcal{X} \quad u \in \mathcal{U}$$

policy improvement: $\mu^{k+1}(x) = \underset{u}{\operatorname{argmin}} Q^{\mu^k}(x, u)$

OBS: the fact that

VI-PI:

$$Q^{\mu^k}(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) Q^{\mu^k}(i, \mu^k(i)) \quad (*)$$

always admit one and only one solution can be also seen by noting that

$$Q^{\mu^k}(x, \mu^k(x)) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) Q^{\mu^k}(i, \mu^k(i)) \quad \text{is the}$$

Bellman equation for value functions, which admits one and only one solution. Hence, (*) universally determines $Q^{\mu^k}(x, \mu^k(x))$ and

$$Q^{\mu^k}(x, \mu^k(x)) = V^{\mu^k}(x) \quad (\text{solution to}). \quad \text{However, once } Q^{\mu^k}(x, \mu^k(x))$$

are determined, (*) determines $Q^{\mu^k}(x, u)$ for the other values of u since right-hand side of (*) becomes given.

$$\text{At every policy evolution: } Q^{\mu^k}(x, u) : Q^{\mu^k}(x, \mu^k(x)) = V^{\mu^k}(x)$$

$$\text{and therefore } Q^{\mu^k}(x, u) = \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) V^{\mu^k}(i)$$

Hence, at every policy improvement

$$\mu^{k+1}(x) = \underset{u}{\operatorname{argmin}} \mathbb{E}[g(x, u, w)] + \gamma \sum_{i=1}^M P(i|x, u) V^{\mu^k}(i)$$

same as PI for value iteration

PI for Q-factor and PI for value function give the same policy at every iteration

convergence and other properties inherited