

Úkol 7: Chceme prát MDP, chceme vytvořit policy.

- 1) V init se jednou vypočítá policy, pak už se znovu nepočítá.
- 2) Chceš to počítat přes $U(s)$ „utility“ a optimální policy.

Policy iteration je těžší implementovat, ale počítá rychleji.

$\hookrightarrow$ čas je důležitý v ReCodExu

Value iteration mi na čas už v náletcích neprojde.

V skutečnosti stačí implementovat precompute-probability-policy-trivial

$$\pi(s) = \operatorname{argmax}_{a \in A} [P(s'/s, a) \cdot U(s')]$$

$$U(s) = R(s) + \max_{a \in A} P(s'/s, a) \cdot U(s')$$

V prezentaci!

Value iteration

inicializace $U(s) = R(s)$

Bellman update:

$$U_{i+1}(s) = R(s) + \gamma \max_{a \in A} P(s'/s, a) \cdot U_i(s')$$

γ — discount factor

Co vlastně chci počítat?

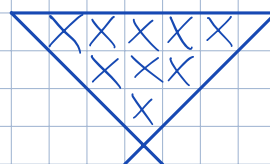
$\rightarrow$ úplně všechno zapomenout vedlejší stavy.
Tedy $\gamma = 1$

$$R(s) = M[s] + \max_{a \in \{L, R, U, D\}} P(s'/s, a) \cdot U(s')$$

Rewards

0 0 0 +1
0 0 0 -1
0 0 0 0

$\gamma = 0.9$ $0.1 \leftarrow \uparrow \rightarrow 0.1$



$$U_1(0_1) = 0 + 0.9 \cdot \max \left(\begin{matrix} \leftarrow & \uparrow & \rightarrow & \downarrow \\ 0, & 0.1 \cdot 1, & 0.8 \cdot 1, & 0.1 \cdot 1 \end{matrix} \right) = 0.9 \cdot 0.8$$

$$U_2(0_2) = 0 + 0.9 \cdot \max \left(0, 0.1 \cdot 0.72, 0.8 \cdot 0.72, 0.1 \cdot 0.72 \right) = \dots$$

$$U_2(1,2) = 0 + 0,9 \cdot \max \left(0,1 \cdot 0,72, \underbrace{0,8 \cdot 0,72 + 0,1 \cdot (-1)}, \dots \right)$$

Jdu přes všechny možnosti
obdobně

Jak zjistit správnou akci? $\pi(s) = \underset{a \in A}{\operatorname{argmax}} \left[\sum_{s'} P(s'/s, a) \cdot U(s') \right]$

Policy iteration

Inicializuj náhodnou policy $\pi_0(s)$

do

evaluate policy

$$U(s) = R(s) + \gamma \cdot \sum_{s'} P(s'/s, \pi(s)) \cdot U(s')$$

$$\text{if } \max_{a \in A} \sum_{s'} P(s'/s, a) \cdot U(s) > \sum_{s'} P(s'/s, \pi(s)) \cdot U(s)$$

potom existuje lepší akce než v policy, update policy