

Sistemi Intelligenti Reinforcement Learning: Q-learning

Alberto Borghese

Università degli Studi di Milano
Laboratorio di Sistemi Intelligenti Applicati (AIS-Lab)
Dipartimento di Informatica

alberto.borghese@unimi.it

Barto and Sutton, 4.7, 6.4, 6.5



A.A. 2025-2026

1/33

<http://borghese.di.unimi.it/>

1



Sommario



Q-learning

Esempi

A.A. 2025-2026

2/33

<http://borghese.di.unimi.it/>

2

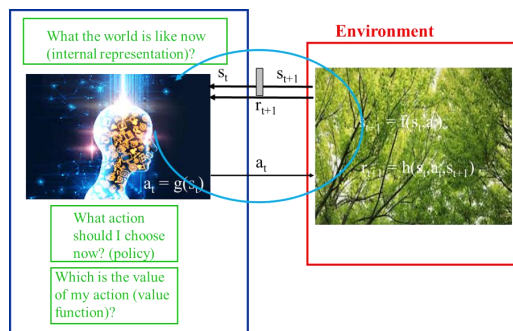


Value Function?

La Value Function deriva:

- dalla visione della Programmazione Dinamica e dell'ottimizzazione.
- dalla psicologia sperimentale: rinforzo totale a lungo termine.

Ma è proprio necessario conoscere esattamente la Value function? In fondo a noi interessa determinare la Policy.



A.A. 2025-2026

3/33

3

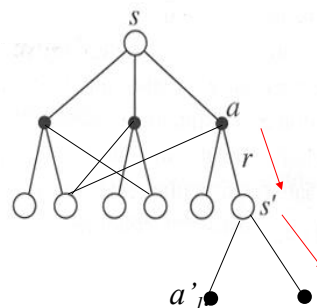


La policy in SARSA

$$Q_{k+1}^{\pi}(s, a) = Q_k^{\pi}(s, a) + \alpha[r' + \gamma Q_k^{\pi}(s', a') - Q_k^{\pi}(s, a)]$$

- 1) Apprendiamo il valore di $Q^{\pi}(s, a)$, $\forall s \forall a$, per una policy data (**on-policy**).
- 2) Dopo avere appreso la funzione $Q^{\pi}(s, a)$, possiamo **modificare la policy, $\pi'(s, a)$** , in modo da migliorarla (**policy improvement**)
- 3) Dopo avere modificato la policy devo apprendere la nuova $Q^{\pi'}(s, a)$

s = state, a = action, r = reward, s = state, a = action

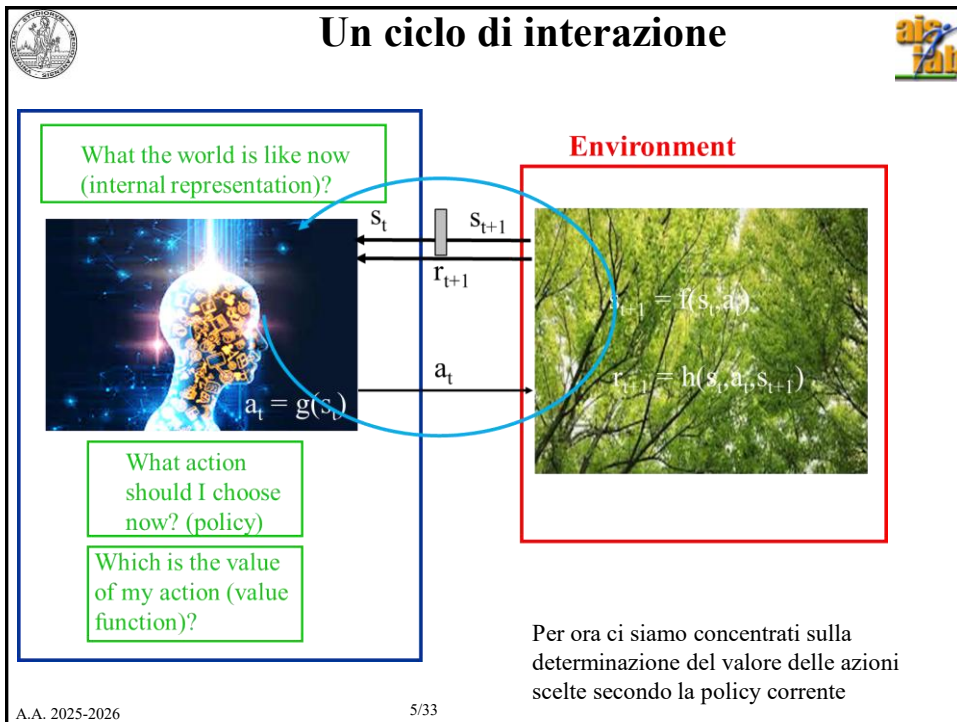


On-policy learning

A.A. 2025-2026

4/33

4



5

Value iteration

Iterative policy evaluation

$$Q_{k+1}^\pi(s_t, a_t) = \sum_{s'} P_{s \rightarrow s' | a} \left\{ R_{s, s', a} + \gamma \sum_{a'} \pi(s', a') Q_k^\pi(s'_{t+1}, a'_{t+1}) \right\}$$

Invece di considerare una policy stocastica, consideriamo l'azione migliore:

$$Q_{k+1}(s_t, a_t) = \max_{a'} \sum_{s'} P_{s \rightarrow s' | a} \left\{ R_{s, s', a} + \gamma \sum_{a'} \pi(s', a') Q_k(s', a') \right\}$$

$\forall s$

A.A. 2025-2026 6/33

6



Off-policy Temporal Difference: Q-learning

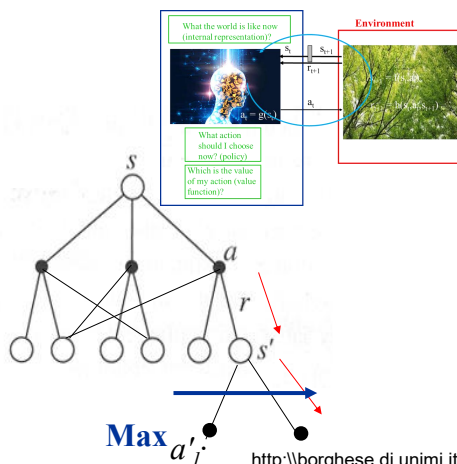


$$Q_{k+1}(s, a) = Q_k^\pi(s, a) + \alpha \left[r' + \gamma \max_{a'} Q_k(s', a') - Q_k^\pi(s, a) \right]$$

Non imparo semplicemente la funzione valore $Q^\pi(\cdot)$, ma la funzione valore $Q^*(\cdot)$ ottima.

In s , scelgo un ramo del grafo, e poi **decido** come continuare, guardando il reward a lungo termine stimato per le diverse azioni in quel passo di iterazione.

Eventualmente cambio subito policy, $a' = \pi(s') \rightarrow a'_{\text{new}} = \pi^*(s')$ senza aspettare di avere stimato esattamente $Q^\pi(s, a)$.



A.A. 2025-2026

7/33

7



Q-learning algorithm (progetto)



```

Q(s,a) = 0;           // ∀ s, ∀ a,
Policy data;         // deterministica o stocastica
Repeat               // for each episode
{ s = s0; α = α*reduction_factor; // decremento il coefficiente di aggiornamento α
  Policy_stable = .true.;
  Repeat             // for each step of the single episode
  { a = Policy(s);   // Policy
    s_next = NextState(s,a); // non nota all'agente
    reward = Reward(s, s_next, a); // non nota all'agente
    a_next_pol = Policy(s_next); // azione in s_next secondo policy nota ad agente
    a_next = argmax_a (Q(s_next, a)); // azione greedy in s_next

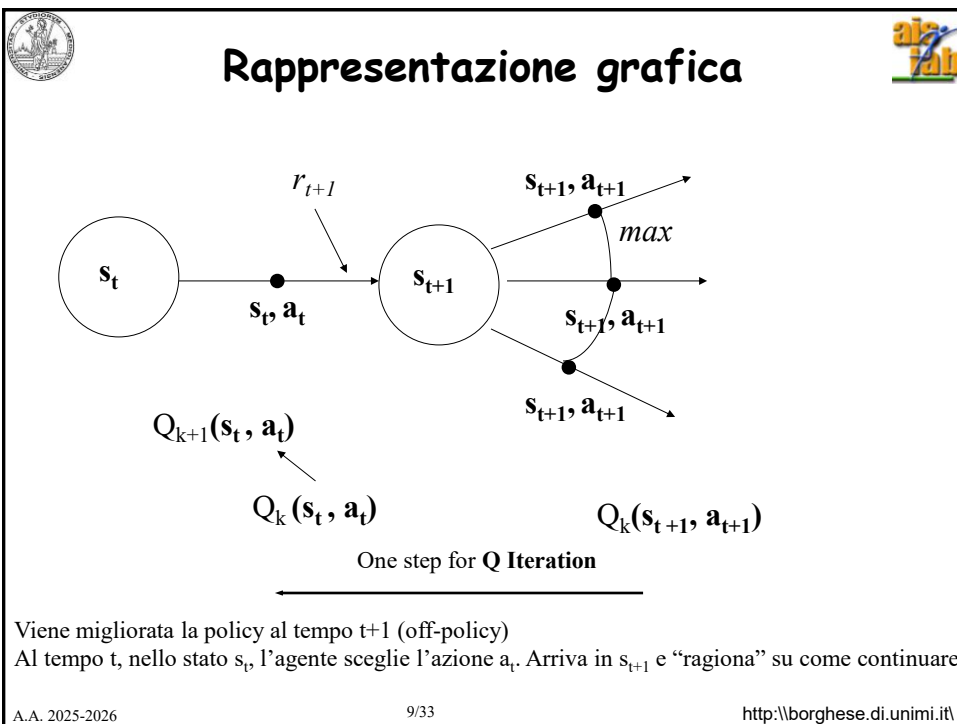
    if (a_next_pol != a_next) // se esiste un'azione a' migliore
    { UpdatePolicy(s_next, a_next); // scelgo a_next in s_next da qui in poi:
      Policy_stable = .false.; // Change in the policy
    }
    endif; // Aggiorno la policy
    Q(s,a) = Q(s,a) + α [reward + γ Q(s_next, a_next) - Q(s,a)]; // aggiorno Q(s,a)
    s = s_next;
  } // until last state
} // until the end of learning (convergence of Q(s,a) to true Q(s,a) ∀ s, ∀ a, for policy π(s,a) )
We may decide to quit when the policy does not change for several episodes.
  
```

A.A. 2025-2026

8/33

<http://borghese.di.unimi.it/>

8



9

Osservazioni

$\pi(s,a)$: l'agente sceglie l'azione ottima

$Q(s, a)$ converge al valore vero (della policy ottima)

Nella pratica la convergenza non viene principalmente valutata sulle variazioni di Q ,
 ma sulla stabilità della policy identificata.

$$Q_{k+1}(s, a) = Q_k(s, a) + \alpha \left[r' + \gamma \max_{a'} Q_k(s', a') - Q_k(s, a) \right]$$

L'operazione di $\max$ può essere un "hard" max o un "soft" max. Si possono considerare policy ϵ -greedy.

Q-learning è **off-policy** perchè la policy viene variata all'interno dell'algoritmo, non rimane la stessa durante tutto l'apprendimento.

A.A. 2025-2026 10/33

10



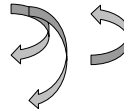
Meccanismo di apprendimento nel RL



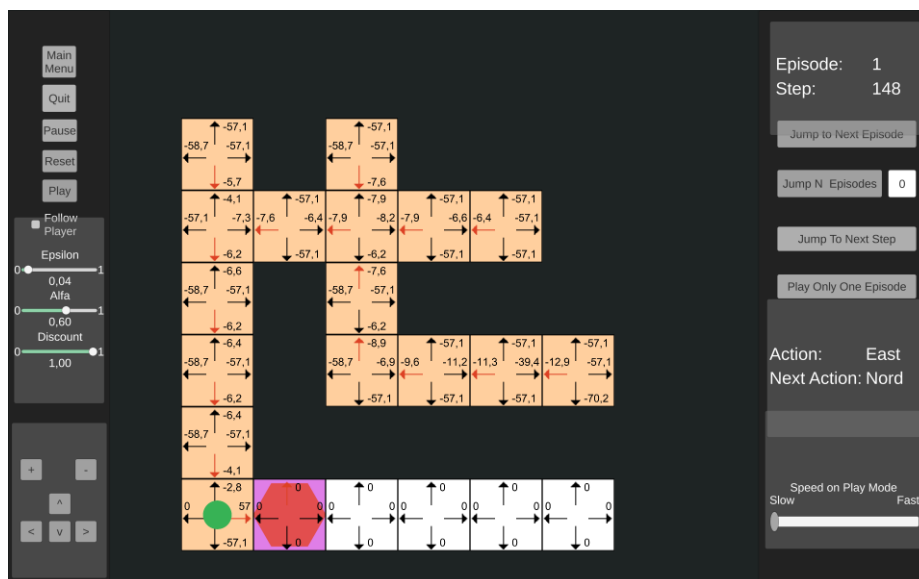
Inizializzazione: se l'agente non agisce sull'ambiente non succede nulla. Occorre specificare una policy iniziale.

Ciclo dell'agente (le tre fasi sono sequenziali):

- 1) Implemento una policy ($\pi(s,a)$)
- 2) Calcolo la Value function ($Q^\pi(s,a)$)
- 3) Aggiorno la policy



Labirinto (S. Abate)





Sommario



Q-learning

Esempi

A.A. 2025-2026

13/33

<http://borghese.di.unimi.it/>

13



Example 1 - Q Learning Update



Esempio tratto dai lucidi del corso di Brian C. Williams su RL.

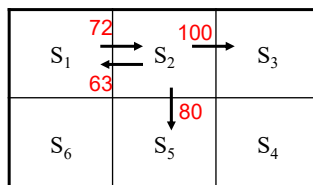
Modificati dalle slide di: Manuela Veloso, Reid Simmons, & Tom Mitchell, CMU

6 stati $\{s_1 \dots s_6\}$

Azioni: {su, destra, giù, sinistra}

Reward istantaneo = 0

Inizializzo $Q(s,a)$ – in rosso.



In rosso i valori di $Q(s,a)$.

Nessun reward istantaneo:
 $r(s,a,s') = 0$.

A.A. 2025-2026

14/33

<http://borghese.di.unimi.it/>

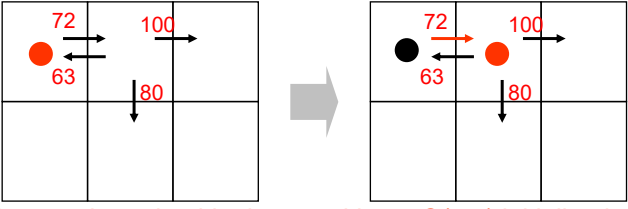
14



Example 1 - Q Learning Update



$\gamma = 0.9$
 $S_{ini} = S_1$



0 reward received in the transitions. $Q(s,a)$ initialized $\neq 0$

S_1	S_2	S_3
S_6	S_5	S_4

In rosso i valori di $Q(s,a)$.
Nessun reward istantaneo.

Apprendimento della funzione valore Q. Versione Q-learning.

Iniziamo con una policy ben definita. Supponiamo:

- $a(s_1) = \text{right}$
- $a(s_2) = \text{down}$
- > $Q_{k+1}(s_1, dx) = ?$

A.A. 2025-2026
15/33
<http://borghese.di.unimi.it/>

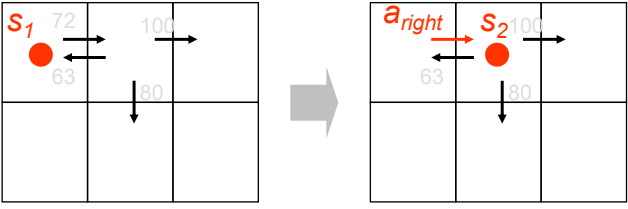
15



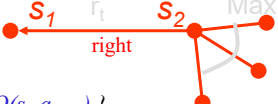
Example 1 - Q Learning: Policy Update



$\gamma = 0.9$
 $\alpha = 0.1$
 $a(s_1) = \text{dx}$
 $a(s_2) = \text{down}$



0 reward received in the transition



$$Q(s_1, a_{right}) = Q(s_1, a_{right}) + \alpha \{ r(s_1, a_{right}, s_2) + \gamma \max_a \{ Q(s_2, a) \} - Q(s_1, a_{right}) \}$$

$$= 72 + \alpha \{ 0 + 0.9 \max_a \{ 63, 80, 100 \} - Q(s_1, a_{right}) \}$$

Correzione di $Q(s_1, a_{right})$
 Correzione dell'azione in s_2 da down a right
 Viene modificata la policy da s_2 in poi

$Q(s_2, a_{down}) = 80$
 $Q(s_2, a_{right}) = 100$
 $Q(s_2, a_{left}) = 63$

A.A. 2025-2026
16/33
<http://borghese.di.unimi.it/>

16



Example 1 - Q Learning Update



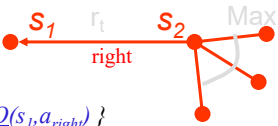
$\gamma = 0.9$
 $\alpha = 0.1$
 $a(s_2) = \text{down}$

s_1	72	100
	63	80

→

a_{right}	s_2	100
	63	80

0 reward received in the transition



$$\begin{aligned}
 Q(s_1, a_{\text{right}}) &= Q(s_1, a_{\text{right}}) + \alpha [r(s_1, a_{\text{right}}, s_2) + \gamma \max_a Q(s_2, a) - Q(s_1, a_{\text{right}})] \\
 &= 72 + \alpha \{0 + 0.9 \max_a \{63, 80, 100\} - 72\} \\
 &= 72 + \alpha (0 + 0.9 \cdot 100 - 72) = 72 + 0.1 \cdot 18 = 73.8 \text{ (da 72 passa a 73.8)}
 \end{aligned}$$

Correzione di $Q(s_1, a_{\text{right}})$

Correzione dell'azione in s_2 da down a right

La correzione di $Q(s_1, a_{\text{right}})$ va a 0 quando

$Q(s_1, a_{\text{right}}) = 90$ (cf. $\alpha = 1$ invece di 0.1)

$Q(s_2, a_{\text{down}}) = 80$

$Q(s_2, a_{\text{right}}) = 100$

$Q(s_2, a_{\text{left}}) = 63$

A.A. 2025-2026
17/33
<http://borghese.di.unimi.it/>

17



Example 1 - Q Learning Update series



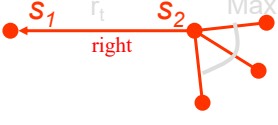
$\gamma = 0.9$
 $\alpha = 0.1$
 $a(s_2) = \text{down}$

s_1	72	100
	63	80

→

a_{right}	s_2	100
	63	80

0 reward received in the transition



$$\begin{aligned}
 Q(s_1, a_{\text{right}}) &= 72 + \alpha (90 - 72) = 72 + 1.8 = 73.8 \quad \text{trial 1} \\
 Q(s_1, a_{\text{right}}) &= 73.8 + \alpha (90 - 73.8) = 73.8 + 1.62 = 75.42 \quad \text{trial 2} \\
 Q(s_1, a_{\text{right}}) &= 75.42 + \alpha (90 - 75.42) = 75.42 + 1.458 = 76.878 \quad \text{trial 3} \\
 Q(s_1, a_{\text{right}}) &= 76.878 + \alpha (90 - 76.878) = 76.878 + 1.3122 = 78.1902 \quad \text{trial 4} \\
 Q(s_1, a_{\text{right}}) &= 78.1902 + \alpha (90 - 78.1902) = 78.1902 + 1.458 = 79.37118 \quad \text{trial 5} \\
 Q(s_1, a_{\text{right}}) &= 79.37118 + \alpha (90 - 79.37118) = 79.37118 + 1.458 = 80.434062 \quad \text{trial 6} \\
 &\dots\dots
 \end{aligned}$$

Si ottiene una serie che converge al valore asintotico 90 (asintoticamente)

A.A. 2025-2026
18/33
<http://borghese.di.unimi.it/>

18



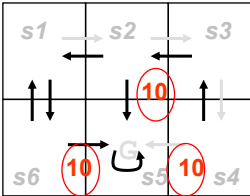
Example 2: Q-Learning Iterations



- Stati: $\{s_1, \dots, s_6\}$
- Azioni: $\{dx, sx, su, giù\}$
- Reward istantaneo solo in alcune transizioni (in rosso e cerchiato).
- $Q(s,a) = 0$ per tutti gli stati.
- Stato iniziale: s_1
- Initial selected policy: move clockwise;

E.g. videogioco.
In G rimango in loop

$\alpha = 1$
 $\gamma = 0.8.$



A.A. 2025-2026

19/33

<http://borghese.di.unimi.it/>



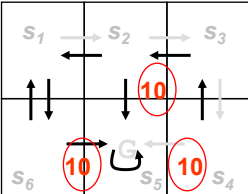
Example 2: Q-Learning Iterations



- Start at upper left; Initial selected policy: move clockwise; $Q(s,a)$ initially 0; $\gamma = 0.8$.
- Reward solo nelle transizioni.

$$Q_{k+1}^{\pi}(s_1, E) = Q_k^{\pi}(s_1, E) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_2, a') - Q_k^{\pi}(s_1, E) \right]$$

Reward istantaneo in rosso e cerchiato



$Q_{k+1}^{\pi}(s_1, E) = 0 + 1[0 + 0.8 \times 0 - 0] = 0$

E.g. videogioco.
In G rimango in G - loop

$Q(s_1, \text{East})$	$Q(s_2, \text{East})$	$Q(s_3, \text{South})$	$Q(s_4, \text{West})$
0			

A.A. 2025-2026

20/33

<http://borghese.di.unimi.it/>



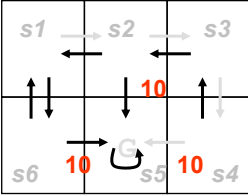
Q-Learning Iterations - trial 1



- Start at upper left – move clockwise; table initially 0; $\gamma = 0.8$; $\alpha = 1$

$$Q_{k+1}^{\pi}(s, a) = Q_k^{\pi}(s, a) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s', a') - Q_k^{\pi}(s, a) \right]$$

$$Q_{k+1}^{\pi}(s_3, S) = 0 + 1[0 + 0.8 \times 0 - 0] = 0$$



$Q(s1, E)$	$Q(s2, E)$	$Q(s3, S)$	$Q(s4, W)$
0	0	0	

A.A. 2025-2026
21/33
<http://borghese.di.unimi.it/>

21



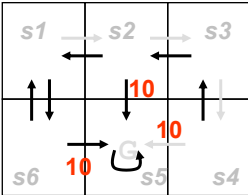
Q-Learning Iterations - trial 1



- Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_4, W) = Q_k^{\pi}(s_4, W) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_3, a') - Q_k^{\pi}(s_4, W) \right]$$

$$Q_{k+1}^{\pi}(s_4, W) = 0 + 1[10 + 0.8 \times 0 - 0] = 10$$



$Q(s1, E)$	$Q(s2, E)$	$Q(s3, S)$	$Q(s4, W)$
0	0	0	10

A.A. 2025-2026
22/33
<http://borghese.di.unimi.it/>

22



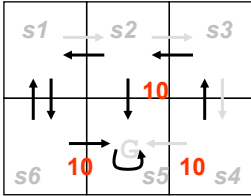
Q-Learning Iterations - trial 2



□ Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_3, S) = Q_k^{\pi}(s_3, S) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_4, a') - Q_k^{\pi}(s_3, S) \right]$$

$$Q_{k+1}^{\pi}(s_3, S) = 0 + 1[0 + 0.8 \{ \max, 10, 0 \} - 0] = 8$$



$Q_k^{\pi}(s_4, W)$

$Q_k^{\pi}(s_4, N)$

Q(s1,E)	Q(s2,E)	Q(s3,S)	Q(s4,W)
0	0	0	10
0	0	8	

A.A. 2025-2026
23/33
<http://borghese.di.unimi.it/>

23



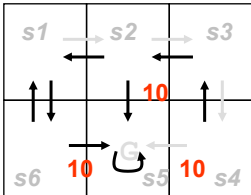
Q-Learning Iterations - trial 2



□ Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_4, W) = Q_k^{\pi}(s_4, W) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_3, a') - Q_k^{\pi}(s_4, W) \right]$$

$$Q_{k+1}^{\pi}(s_4, W) = 10 + 1[10 + 0.8 \times 0 - 10] = 10$$



$Q_k^{\pi}(s_5, \cdot)$ goal

Q(s1,E)	Q(s2,E)	Q(s3,S)	Q(s4,W)
0	0	0	10
0	0	8	10

A.A. 2025-2026
24/33
<http://borghese.di.unimi.it/>

24



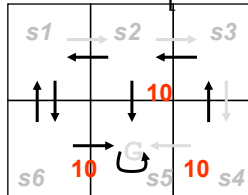
Q-Learning Iterations - trial 3



- Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_2, E) = Q_k^{\pi}(s_2, E) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_3, a') - Q_k^{\pi}(s_2, E) \right]$$

$$Q_{k+1}^{\pi}(s_2, E) = 0 + 1 \left[0 + 0.8 \times \max_{a'} \{8, 0\} - 0 \right] = 6.4$$



$Q_k^{\pi}(s_3, S)$

$Q_k^{\pi}(s_3, W)$

$Q(s1, E)$	$Q(s2, E)$	$Q(s3, S)$	$Q(s4, W)$
0	0	0	10
0	0	8	10
0	6.4		

A.A. 2025-2026

25/33

<http://borghese.di.unimi.it/>

25



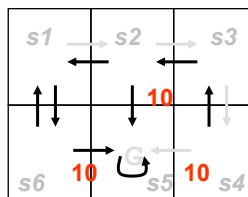
Q-Learning Iterations - trial 3



- Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_3, S) = Q_k^{\pi}(s_3, S) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_4, a') - Q_k^{\pi}(s_3, S) \right]$$

$$Q_{k+1}^{\pi}(s_3, S) = 0 + 1 \left[0 + 0.8 \{ \max, 10, 0 \} - 0 \right] = 8$$



$Q_k^{\pi}(s_4, W)$

$Q_k^{\pi}(s_4, N)$

$Q(s1, E)$	$Q(s2, E)$	$Q(s3, S)$	$Q(s4, W)$
0	0	0	10
0	0	8	10
0	6.4	8	10

A.A. 2025-2026

26/33

<http://borghese.di.unimi.it/>

26



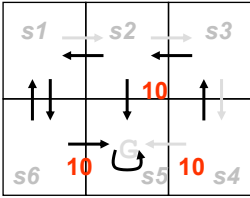
Q-Learning Iterations - trial 4



Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_1, E) = Q_k^{\pi}(s_1, E) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_2, a') - Q_k^{\pi}(s_1, E) \right]$$

$$Q_{k+1}^{\pi}(s_1, E) = 0 + 1 \left[0 + 0.8 \times \max_{a'} \{6.4, 0, 0\} - 0 \right] = 5.12$$



$Q_k^{\pi}(s_2, W)$
 $Q_k^{\pi}(s_2, E)$
 $Q_k^{\pi}(s_2, S)$

$Q(s_1, E)$	$Q(s_2, E)$	$Q(s_3, S)$	$Q(s_4, W)$
0	0	0	10
0	0	8	10
0	6.4	8	10
5.12	6.4	8	10

Potrei migliorare la policy: dovrei scegliere South in s_2 . **Come?** <http://borghese.di.unimi.it/>

27



Q-Learning Iterations: improving policy

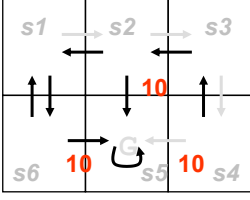


Start at upper left – move clockwise; $\gamma = 0.8$; $\alpha = 1$.

Mossa ϵ -greedy in s_2 (invece che a = Est, scelgo a = South, cambio azione)

$$Q_{k+1}^{\pi}(s_2, S) = Q_k^{\pi}(s_2, S) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_5, a') - Q_k^{\pi}(s_2, S) \right]$$

$$Q_{k+1}^{\pi}(s_2, S) = 0 + 1 [10 + 0.8 \times 0 - 0] = 10$$



$Q_k^{\pi}(s_5, \cdot)$

$Q(s_1, E)$	$Q(s_2, E)$	$Q(s_2, S)$	$Q(s_3, S)$	$Q(s_4, W)$
0	0	0	0	10
0	0	0	8	10
0	6.4	0	8	10
5.12	6.4	10	8	10

28



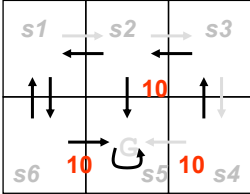
Q-Learning Iterations: policy changed!



Start at upper left – move clockwise; $\gamma = 0.8$

$$Q_{k+1}^{\pi}(s_1, E) = Q_k^{\pi}(s_1, E) + \alpha \left[r' + \gamma \max_{a'} Q_k^{\pi}(s_2, a') - Q_k^{\pi}(s_1, E) \right]$$

$$Q_{k+1}^{\pi}(s_1, E) = 5.12 + 1 \left[0 + 0.8 \times \max_{a'} \{6.4, 10, 0\} - 5.12 \right] = 8$$



$Q_k^{\pi}(s_2, E)$ (blue arrow to 6.4)
 $Q_k^{\pi}(s_2, S)$ (red arrow to 10)
 $Q_k^{\pi}(s_2, W)$ (blue arrow to 0)

$Q(s1, E)$	$Q(s2, E)$	$Q(s2, S)$	$Q(s3, S)$	$Q(s4, W)$
0	0	0	0	10
0	0	0	8	10
0	6.4	0	8	10
8	6.4	10	8	10



Proprietà del rinforzo



L'ambiente o l'interazione può essere complessa.

Il rinforzo può avvenire solo dopo una più o meno lunga sequenza di azioni (delayed reward).

E.g. agente = giocatore di scacchi.
 ambiente = avversario.

Problemi collegati:

- temporal credit assignment.
- structural credit assignment.

L'apprendimento non è più da esempi, ma dall'osservazione del proprio comportamento nell'ambiente.

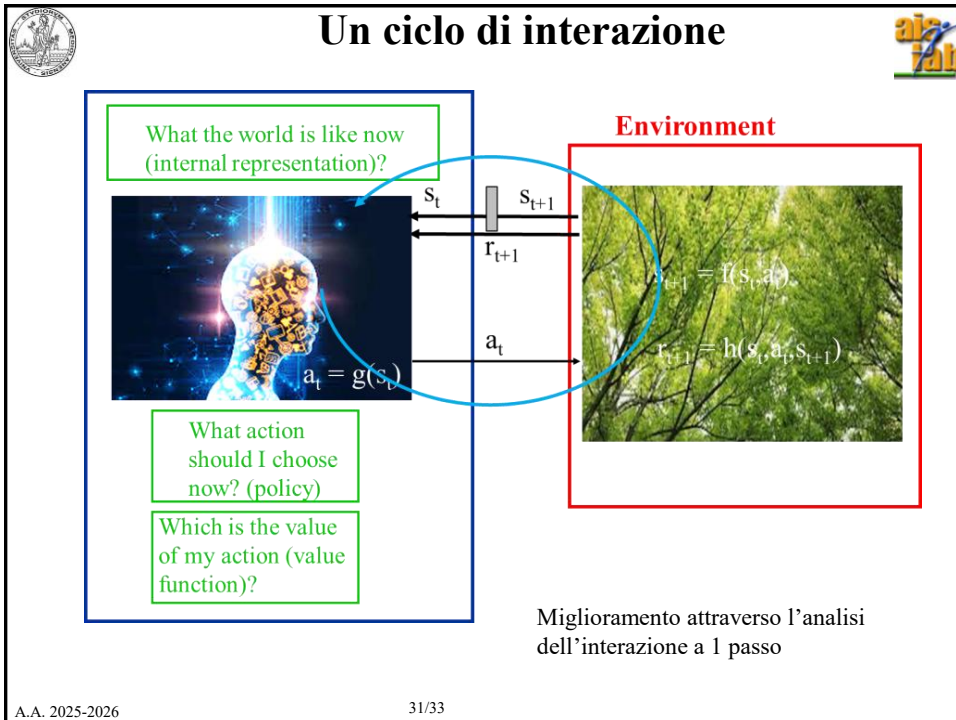
Utilizzo delle equazioni di Bellman

Utilizzo una "porzione" (sample) di esperienza per migliorare la policy.

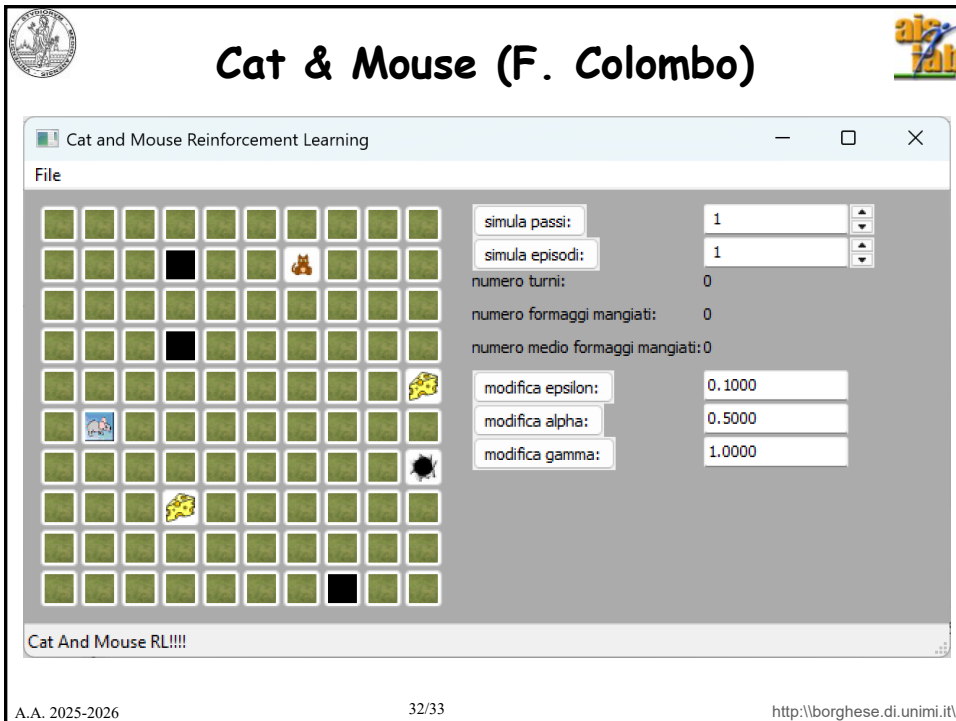
A.A. 2025-2026

30/33

<http://borghese.di.unimi.it/>



31



32



Sommario



Q-learning

Esempi