

Reinforcement Learning

Grundlagen Grundbegriffe Grundübungen

Alfatraining Kurs 2025-08 Dozent: Karsten Flügge

Themen des Tages

Tag 6

Deep Reinforcement Learning

TD & Monte-Carlo

Temporal-Difference (TD)

Unsere Quality Update Methode ist ein Beispiel von **Temporal-Difference (TD) lernen**, da bei jedem Update die **Differenz** zwischen zwei Schätzungen verwendet wird:

QUALITY UPDATE

$$q(a') = q(a | s) + [r + \gamma \max_a q(s', a) - q(a | s)]$$

$$q'(a | s) = q(a | s) + \alpha * [r + \gamma \max_a q(s', a) - q(a | s)]$$

α learning rate: wie viel wollen wir unsere Schätzung updaten

Q-Learning (off-policy) : wende die Formel als pre-training an.

SARSA (on-policy) : wende die selbe Formel mehrfach pro episode (oder per replay) an:

$$q_target = r + \gamma \max_a q(s', a')$$

$$q'(a) += \alpha * [q_target - q(a | s)]$$

Temporal-Difference (TD) learning bezeichnet eine ganze Klasse von ähnlichen Ansätzen, in welchen in der Update Formel die zeitliche **Differenz zwischen zwei Schätzungen** vorkommt.

All Q-learning is TD learning, but not all TD learning is Q-learning.

Temporal-Difference (TD)

QUALITY UPDATE

$\text{best_reward} := \max_a q(s', a')$

$\text{target} := \text{current_reward} + \gamma * \text{best_reward}$

$\text{target} = r + \gamma \max_a q(s', a')$

$\text{estimate} := q(a | s)$

$q'(a | s) = q(a | s) + \alpha * [r + \gamma \max_a q(s', a') - q(a | s)]$

$q'(a | s) = q(a | s) + \alpha * [\text{target} - \text{estimate}]$

α learning rate: wie stark viel wollen wir unsere Schätzung updaten?

TD = Ein Zeitschritt $(s, a) \Rightarrow (s', a')$

SARSA = state action reward state action

