

Monte-Carlo & Temporal-Difference Learning

Grundlagen & Wiederholung

Ein Agent lebt in seiner Welt, macht Aktionen, bekommt Belohnungen – und will lernen, wie er sich verhalten soll, um langfristig möglichst viel Belohnung zu sammeln.

Die drei zentralen Fragen:

1. *Wie gut ist es, in Zustand S zu sein?* → **Value Function V(s)**
2. *Wie gut ist es, in Zustand S die Aktion A zu wählen?* → **Action-Value Function Q(s,a)**
3. *Welche Aktion soll ich wählen?* → **Policy π**

Bellman-Gleichung (intuitiv):

„Der Wert eines Zustands = Sofortige Belohnung + Wert des nächsten Zustands (abgezinst)“

Missverständnisse ⚠

Missverständnis	Richtigstellung
„Belohnung = Wert“	Nein! Belohnung ist sofortig , Wert ist langfristig
„ $\gamma=1$ ist am besten“	Kann zu Instabilität führen, kein Anreiz für schnelle Lösung
„Policy = deterministisch“	Kann stochastisch sein: $\pi(a s) \in [0,1]$

Frage: Warum brauchen wir einen Discount-Faktor $\gamma < 1$?

- a) Damit der Agent schneller lernt
- b) Damit zukünftige Belohnungen weniger zählen als sofortige
- c) Damit die Policy deterministisch wird

Monte-Carlo Methoden

Erklärung (intuitiv):

Kernidee: Lernen durch vollständige Erfahrung.

🎯 „Spiele das Spiel bis zum Ende, dann schau was insgesamt rausgekommen ist.“

Monte-Carlo Prinzip:

1. Führe eine **komplette Episode** durch ($S_0, A_0, R_1, S_1, \dots, S_t$)
2. Berechne den **tatsächlichen Return** rückwärts
3. Update den Schätzwert V(s) in Richtung des tatsächlichen Returns

Return berechnen:

$$G_t = R_{t+1} + \gamma \cdot R_{t+2} + \gamma^2 \cdot R_{t+3} + \dots + \gamma^{T-t} \cdot R_t$$

Update-Regel:

$$V(s) \leftarrow V(s) + \alpha \cdot [G - V(s)]$$

↑ ↑
Lernrate "Überraschung"

Frage: Warum hat MC eine hohe Varianz?

Frage: Kann man MC für ein Schachspiel einsetzen?

Frage: Was ist α (alpha) in der Update-Regel?

- a) Der Discount-Faktor
- b) Die Lernrate (Schrittgröße)
- c) Die Explorations-Rate

Beispiel:

Schachbrett / Gridworld

[][][][👤]

[][🔴][][][]

[][][][]

[👤][][][💀]

Agent (👤) → Ziel (👤), vermeide Tod (💀)

Episode 1: 👤→→→↑→👤
Belohnungen: 0, 0, 0, 0, +10
G (rückwärts berechnet):

Zeitschritt:	t=0	t=1	t=2	t=3	t=4
Reward:	0	0	0	0	+10
Return G:	7.29	8.1	9.0	10.0	10.0

← (γ=0.9)

NACH der Episode: Update aller besuchten Zustände!

First-Visit vs. Every-Visit MC

Episode: $S_1 \rightarrow S_2 \rightarrow S_1 \rightarrow S_3 \rightarrow \text{Ende}$
 ↑
 S_1 wird zweimal besucht!

First-Visit MC: Nur 1. Besuch von S_1 zählt
Every-Visit MC: Beide Besuche von S_1 zählen

✅ Vorteile	❌ Nachteile
Kein Bias – lernt aus echten Returns	Braucht komplette Episoden
Einfach zu verstehen	Hohe Varianz
Kein Modell der Umgebung nötig	Kein Online-Learning möglich
Unabhängig von Initialwerten	Nur für episodische Tasks

Missverständnisse ⚠

Missverständnis	Richtigstellung
„MC ist langsam, also schlecht“	Für viele echte Probleme (z.B. Brettspiele) sehr effektiv
„First-Visit ist besser als Every-Visit“	Beide konvergieren – Kontext entscheidet
„MC lernt aus jedem Schritt“	Nein! Nur nach kompletter Episode

Überblick

Lernt aus:		
	Komplette Episode	Jeden Schritt
On-Policy	MC (on-policy)	SARSA Expected SARSA
Off-Policy	MC (off-policy)	Q-Learning

Eigenschaft	MC	TD(0)	SARSA	Exp. SARSA	Q-Learning
Bootstrapping	✗	✓	✓	✓	✓
Online Learning	✗	✓	✓	✓	✓
On/Off Policy	Beide	–	On	On	Off
Varianz	Hoch	Mittel	Mittel	Niedrig	Mittel
Bias	Kein	Mittel	Mittel	Mittel	Mittel
Konvergenz-Ziel	V(s)	V(s)	Q*	Q*	Q*

Entscheidungsbaum:

Ist die Task episodisch?

- Nein → TD-Methoden (SARSA, Q-Learning)
- Ja → Brauche ich schnelles Online-Learning?
 - Nein (und viele Daten) → Monte-Carlo
 - Ja → Muss ich konservativ/sicher sein?
 - Ja (z.B. Roboter, echte Klippen) → SARSA
 - Nein → Habe ich viele Aktionen?
 - Ja → Expected SARSA (geringere Varianz)
 - Nein → Q-Learning (einfach, konvergiert zu Q*)

