

מערכות דינאמיות ובקרה

לביא שפיגלמן

מבוא ל-Reinforcement Learning

• הקשר בין DP ו-RL

הקשר בין DP ו-RL

כזכור, בבעיית הבקרה הדיסקרטית, משוואת הדינאמיקה היא

$$x_{t+1} = f(x_t, u_t, t)$$

הcost to go, (שסומן כ- J^0 או V^0) על המסלול האופטימלי הוא

$$V^0(x_k, k) = \min_{\mathbf{u}} \left\{ \phi(x_N, N) + \sum_{t=k}^{N-1} L(x_t, u_t, t) \right\}$$

או, בנוסחת נסיגה

$$V^0(x_k, k) = \min_{u_k} \{ L(x_k, u_k, k) + V^0(f(x_k, u_k, k), k+1) \}$$

ב-DP משתמשים בנוסחא הנ"ל (ובפונקציית המחיר הרגעי, L ומשוואות הדינאמיקה, f) למציאת $V^0(x_k, k)$ לכל x ו- k ע"י איטרציה. כפי שצויין, זמן החישוב, (עבור מספר צעדים, N , הוא $\mathcal{O}(N|\mathcal{U}||\mathcal{X}|)$) עלול להיות גדול ופונקציות הדינאמיקה והמחיר הרגעי לא תמיד ידועות.

התחום, Reinforcement Learning (RL) בא להתמודד עם הבעיות הנ"ל. לשם כך יש צורך בהתאמת מושגים קלה. הloss הרגעי נלקח כשלילי ונקרא reinforcement. מחיר הנקודה הסופית נלקח גם הוא כשלילי ונקרא גם כן reinforcement.

כשם שבבעיית בקרה האופטימלית רצינו למזער את המחיר הכולל, ב-RL רוצים למקסם את (תוחלת) סכום ה-reinforcements.

סימונים (ותרגומם):

שם בבלקרה	סימון בבלקרה	שם בRL
state	$\mathbf{x}$	state
control sig.	$\mathbf{u}$	action
control function	$\mathbf{u}(\mathbf{x}, t)$	(stochastic) policy
state dynamics	$\mathbf{f}(\mathbf{x}, \mathbf{u}, t, \omega)$	(stochastic) state transition
instantaneous loss	$-L$	(stochastic) reward
cost to go	$-J, -V$	commulative expected reward

התמונה הטיפוסית בRL היא של סוכן (agent) שחש את סביבתו באופן חלקי ומייצג את ההערכה שלו של מצבו בעולם ע"י s . בכל צעד, הסוכן מבצע פעולה באופן סטוכסטי (או דטרמיניסטי). התפלגות הפעולות מיוצגת ע"י $Pr(a|s) = \pi(s, a)$ (במונחים של בקרה, $\mathbf{u}(\mathbf{x})$ זו פונקציה הסתברותית של המצב). לאחר ביצוע הפעולה המצב משתנה ע"פ התפלגות (ידועה או לא ידועה)

$$P_{ss'}^a = Pr\{s_{t+1} = s' | s_t = s, a_t = a\}$$

במונחים של בקרה $P_{xx'}^u$ מקבילה להסתברות ש $x' = f(x, u, \omega)$ (הסטוכסטיקות נובעת מרעש תהליך ω). לאחר מכן מתקבל reward (loss במונחים של בקרה). הreward ניתן מתוך התפלגות ותוחלתו

$$R_{ss'}^a = E_{r_{t+1}}\{r_{t+1} | s_t = s, a_t = a, s_{t+1} = s'\}$$

מטרת הסוכן היא לבצע policy אופטימלי, כלומר לבצע פעולות כך ששכום הrewards יהיה מירבי. שכום זה נלקח בדרך כלל כשכום אינסופי מונחת (discounted). השכום של ריצה ספציפית יסומן כך:

$$\begin{aligned} R_t &= r_{t+1} + \gamma r_{t+2} + \gamma^2 r_{t+3} + \dots \\ &= r_{t+1} + \gamma R_{t+1} \end{aligned}$$

ותוחלת השכום (הרווח) תסומן

$$\begin{aligned} V^\pi(s) &= E_\pi\{R_t | s_t = s\} \\ &= E_\pi\{r_{t+1} + \gamma V(s_{t+1}) | s_t = s\} \\ &= \sum_a \pi(s, a) \sum_{s'} P_{ss'}^a [R_{ss'}^a + \gamma V^\pi(s')] \end{aligned} \quad (1)$$

זה כמובן cost to go (בהצגה סטוכסטית), $J(x, t)$ נשים לב כי כאן מניחים שההתפלגות גויות אינן תלויות בזמן (הנחת סטציונריות) והמחיר הוא עבור זמן ריצה אינסופי (תוך דעיכת הrewards ע"י גורם γ) ולכן הפונקציה היא של המצב בלבד. במקרה זה פתרון הDP מצריך מילוי של וקטור ערכים, v , על סמך ערכו הקודם במקום מילוי של עמודות עוקבות במטריצה.

המחיר תחת בקרה אופטימלית, $J^0(x, t)$, יסומן כ

$$V^*(s) = \max_\pi V^\pi(s)$$

והpolicy האופטימלי הוא $\pi^*(s)$. כפי שהגדרנו את $J^1(x, u, t)$ להיות מחיר ביצוע אות בקרה כלשהו, u ולאחריו ביצוע בקרה אופטימלית, כך בRL מוגדרת פונקציית

הֵן optimal action-value (באופן רקורסיבי):

$$\begin{aligned} Q^*(s, a) &= E \{ r_{t+1} + \gamma V^*(s_{t+1}) \} \\ &= E \left\{ r_{t+1} + \gamma \max_{a'} Q^*(s', a') \right\} \\ &= \sum_{s'} P_{ss'}^a \left[R_{ss'}^a + \gamma \max_{a'} Q^*(s', a') \right] \end{aligned}$$

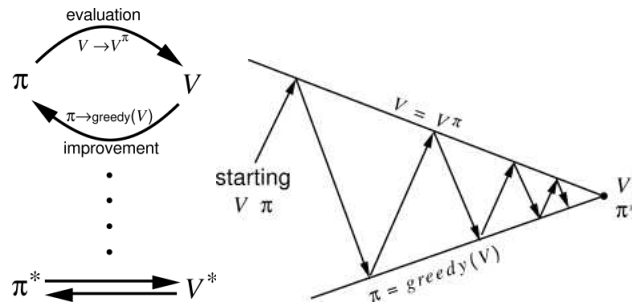
זה האנלוג לתנאי האופטימליות של Bellman מבקרה אופטימלית. הֵן policy האופטימלי במונחים של Q הוא

$$\pi^*(s) = \arg \max_{a \in A(s)} Q^*(s, a)$$

כשם שמצאנו את $J(x, t)$ (ואת $J^1(x, u, t)$) בשיטות DP, ניתן למצוא (בהינתן פונקציות L ו־ f , סטוכסטיות, כלומר בהינתן $R_{ss'}^a$ ו־ $P_{ss'}^a$) policy אופטימלי באותה סיבוכיות זמן ריצה.

כשהדינאמיקה לא ידועה במלואה או המחיר הרגעי אינו ידוע במלואו או זמן הריצה אינו ארוך דיו (כפי שקורה למשל בשש בש בו יש כ- 10^{20} מצבים) RL מציב אלגוריתמים שמתכנסים (בד"כ, תחת תנאים מסוימים אך ריאליים) לפיתרון האופטימלי (כלומר מתכנסים לקיום תנאי האופטימליות של Bellman). ישנה גם הרחבה של RL לזמנים רציפים.

כל שיטות הֵן RL וכן DP (וגם EM) מבצעות איטרציות בהן בצעד אחד ישנו הערכה של הֵן value function עבור הֵן policy הנתונה וצעד שני בו יש שיפור של הֵן policy על סמך הֵן value function המעודכנת.



שיטות מעין זו מכונות Generalized Policy Iteration (GPI) הֵן DP הקלאסי שהוצג כאן מבצע סריקה של כל המצבים ועדכון של $V(s)$ לכל s לפני π מעודכן¹. דבר זה אינו הכרחי. Asynchronous DP מעדכן את $V(s)$ עבור קבוצה חלקית של S (אפילו s בודד) ולאחר מכן מעדכן את הֵן policy, כלומר, במקום לעדכן את וקטור הערכים במלואו על סמך וקטור קודם, מעכענים מצבים בודדים על סמך הוקטור הקודם או, אפילו ערך בודד על סמך שאר הערכים (ללא זיכרון זמני נוסף). באופן זה ניתן לרכז את מאמצי החישוב במצבים שסבירות ההגעה אליהם גבוהה.

¹ בהצגת הנוסחאות שלנו הדבר נבע מכך ש $V_{i,k}$ חושב לכל i ע"י בחירת u שמיזער את V על סמך $V_{:,k+1}$