

知識工学

岡山大学大学院

講師 竹内孔一

本日の内容

■強化学習

使う必要がある場合とは？

■学習データが使えないとき

- 人手による正解が作れない

例)

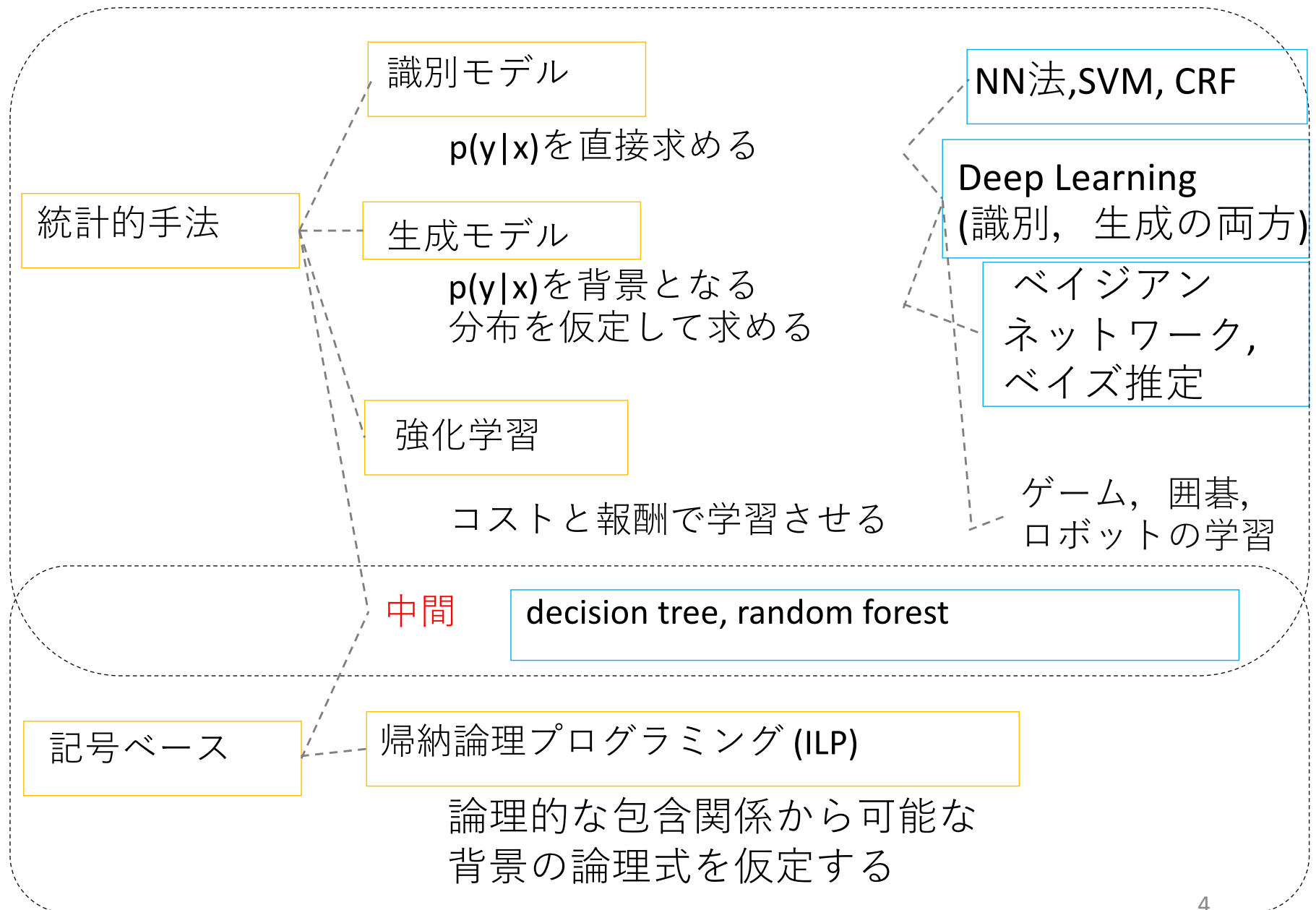
■ロボットの部品を掴む動作/ゲームの動作

- 状況によって腕のどの角度にするかは1通りでは無い
- 1つの正解例(人間が教える)としても違ってても、目標さえ達成すれば腕の角度やスピードに関して幅がある
- 人手で作成する学習データとして可能な全ての組合せでうまくいく例をつくるのは不可能

■強化学習では

- 正解(手順を教える)の代わりに、うまくいったときに報酬を与えることで手順はモデルに生成してもらう

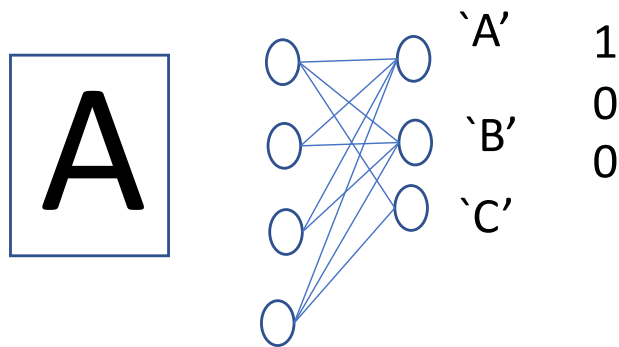
(統計的)パターン認識まわりの学習法



教師あり/教師なし(報酬あり)

教師あり

各データに対して1つの正解
が与えられる (例)文字認識



深層学習など使うと良い

強化学習不要

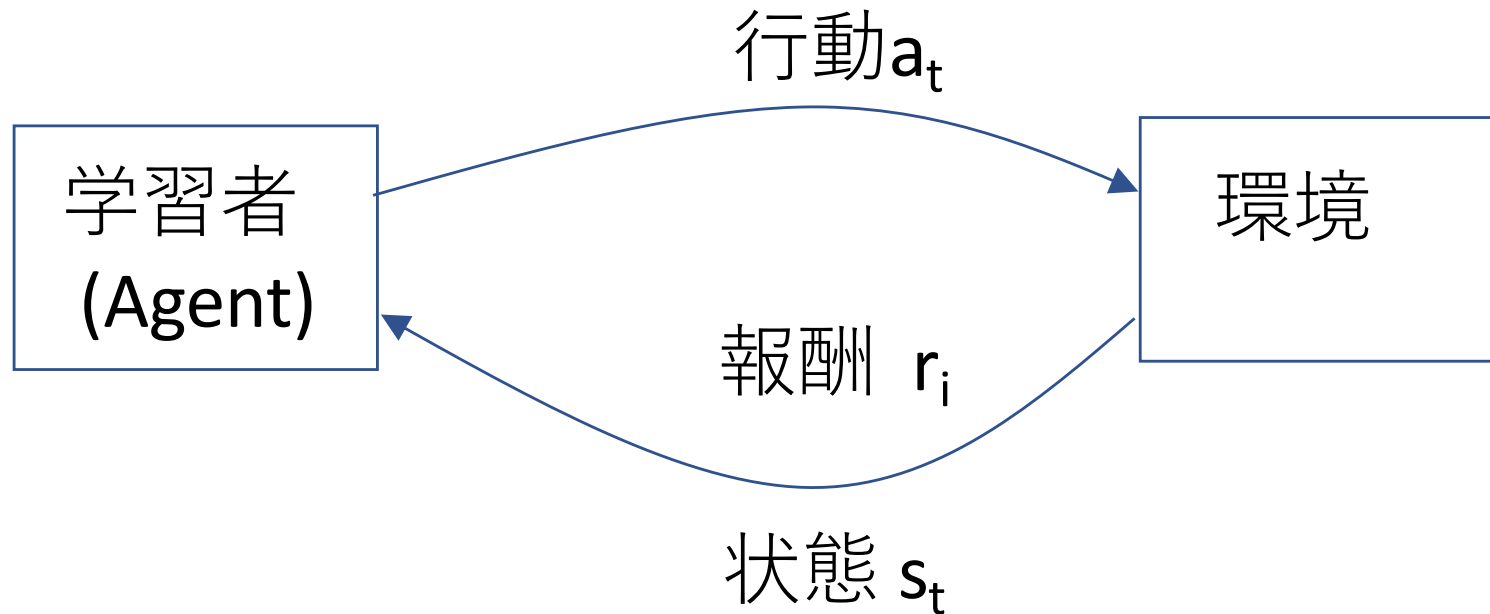
ある状態だけ良い

とり得る値(x)に対して正解は無い
が、時系列後に正解がある
(例)迷路を解く

S	→				
	↓	障			
			障		G

強化学習が使える

強化学習の枠組



学習者は環境から状態 s_t を受け取り、ある行動 a_t をとる。
それにより報酬を得る。

学習者は環境から状態 s_t を受け取り、ある行動 a_t をとる。
それにより環境が s_{t+1} に遷移。報酬 r_{t+1} を得る。

強化学習の枠組

■ 状態 s_t (state)

- エージェントが時刻 t で取る状態

■ 行動 a_t (action)

- エージェントが時刻 t で取る行動

■ 報酬 r_t (reward) ・ 収益 R_t (return)

- 報酬: エージェントが行動により得る値
- 収益: 時刻 t (または $t+1$)以降に得られる報酬の総量

■ 政策 $\pi(s,a)$ (policy)

- 状態 s のときに行動 a をとる関数

■ 価値関数 $V(s)$ (state-value function)

- 状態 s で将来得られる報酬の総量(=収益)の期待値 (つまり予測値)

■ 行動価値関数 $Q(s,a)$ (action-value function)

- 状態 s で行動 a を取るとき将来得られる報酬の総量(=収益)の期待値

■ 環境モデル

- 状態 s で行動 a を取ったとき、次にどういう状態に行くか、報酬はあるかあるとしたらいくらか、エージェントに与える

練習

下記のように時刻 $t=1$ のとき状態 $s_t = (1,1)$ にいるとする. $t=2$ のときに $s_2 = (1,2)$ に移動して, 報酬を2もらったとする. 次に $s_3 = (2,3)$ に移動して報酬1をもらったとする.

問1 s_1 から s_3 に至る行動を示せ. ここでエージェントが取れる行動は上下左右とする

問2 s_2 で得た報酬はいくらか.

問3 s_1 から s_3 まで移動したときの収益はいくらか

問3 次に s_4 に移動したときは報酬がなく, s_5 で報酬3を得た. s_1 から s_5 までの行動と収益を求めよ

	1	2	3
1	s_1	s_2	
2		s_3	
3		s_4	s_5

強化学習の基本枠組(これで全部)

■ マルコフ決定過程

■ 状態遷移モデル

■ エージェントの行動

■ 状態 s_t をで行動 a_t を選択

■ 環境から次状態 s_{t+1} と報酬 r_{t+1} を得る

■ (注) 行動 a_t で報酬 r_{t+1} (教科書などによるので注意)

■ 収益

■ これから得られる報酬の総量 (t は T まで)

■ 将来の収益は割り引いて考える γ (割引率)

■ 状態価値 $V(s)$ vs. 行動価値 $Q(s, a)$

■ (ある状態 s での価値) vs. (状態 s で行動 a の時の価値)

■ 政策 π によって違う値をとる

■ 状態価値 $V^\pi(s)$ vs. 行動価値 $Q^\pi(s, a)$

■ 期待値を求めて、行動選択の指針にする

■ 状態価値を使うか行動価値を使うかはユーザが選択

■ 学習方法 (状態 $V(s)$ か 行動 $Q(s, a)$ の学習か2種)

■ $V(s)$ の学習: TD 学習 (Temporal difference learning)

■ 各状態 s での価値 $V(s)$ が数値として求まる

■ $Q(s, a)$ の学習(1): 方策オン型学習

■ **SARSA** (方策 π に従った学習法)

■ $Q(s, a)$ の学習(2): 方策オフ型学習

■ **Q-learning** (方策は関係無く価値最大の行動をとると固定) 簡単に利用できる

練習

■ 下記の移動問題で各状態 s の価値 $V(s)$ が計算できたとする。政策 π を「価値が最大の状態に進む」とするとき、 S から G に向けてどういう状態遷移をするか状態列を示せ。

■ ここで状態は $(1,1)$ など座標で表すとする

■ 行動は上下左右のみとする

座標	1	2	3
1	S 0.0	5.8	8.1
2	3.5	9.6	G 20.5