

Interactively Shaping Agents via Human Reinforcement

W. Bradley Knox 、 Peter Stone
The University of Texas at Austin
K-CAP-2009

Introduction

- Background: As computational learning agents move into domains that incur real costs (e.g., autonomous driving or financial investment), it will be necessary to learn good policies without numerous high-cost learning trials.
- Goal: learn from human to speed learning

Learning from Advice

Learning from Advice

- suggesting an action when a certain condition is true.
- [1] created a domain-specific natural language interface for giving advice to a reinforcement learner.
- general natural language recognition is unsolved
- Moreover, work still remains on how to embed advice into agents that learn from experience.

[1] G. Kuhlmann, P. Stone, R. Mooney, and J. Shavlik. Guiding a reinforcement learner with natural language advice: Initial results in RoboCup soccer

Learning from demonstration

- Human provide trajectories
- Learn reward or policy

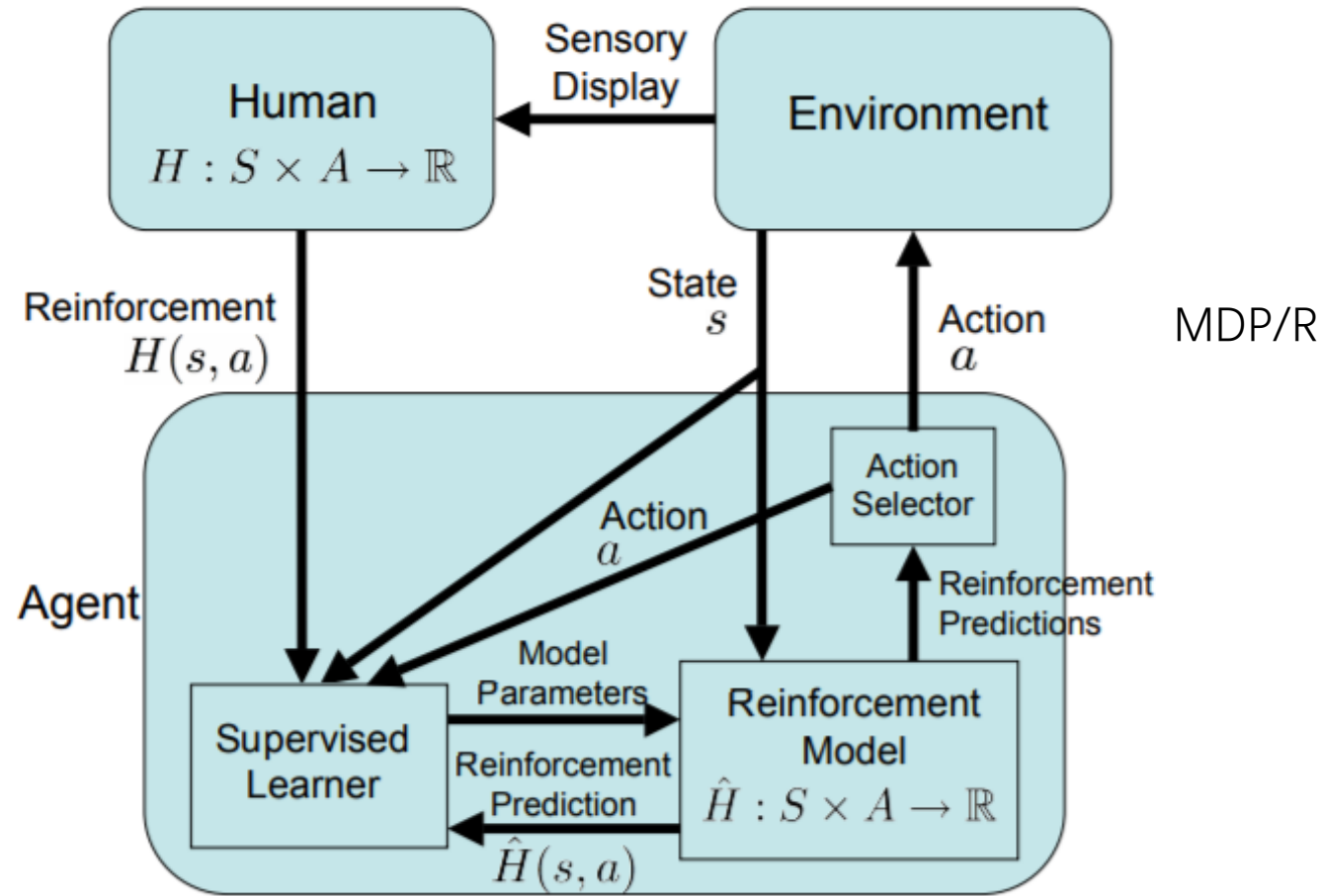
Learning from Reinforcement

Provide human signal to shape the agent

[1] Their agent seeks to maximize the sum of human reinforcement and environmental reward.

A. Thomaz and C. Breazeal. Reinforcement Learning with Human Teachers: Evidence of Feedback and Guidance with Implications for Learning Performance

Framework



Algorithm

Algorithm 1 A general greedy TAMER algorithm

Require: *Input: stepSize*

```
1: ReinfModel.init(stepSize)
2:  $\vec{s} \leftarrow \vec{0}$ 
3:  $\vec{f} \leftarrow \vec{0}$ 
4: while true do
5:    $h \leftarrow \text{getHumanReinfSincePreviousTimeStep}()$ 
6:   if  $h \neq 0$  then
7:      $\text{error} \leftarrow h - \text{ReinfModel.predictReinf}(\vec{f})$ 
8:      $\text{ReinfModel.update}(\vec{f}, \text{error})$ 
9:   end if
10:   $\vec{s} \leftarrow \text{getStateVec}()$ 
11:   $a \leftarrow \text{argmax}_a(\text{ReinfModel.predict}(\text{getFeatures}(\vec{s}, a)))$ 
12:   $\vec{f} \leftarrow \text{getFeatures}(\vec{s}, a)$ 
13:   $\text{takeAction}(a)$ 
14:  wait for next time step
15: end while
```

模型更新

动作选择

Credit assignment

$$\begin{aligned}c_t &= P(\text{event starting at } t_i \text{ was targeted}) \\&= P(\text{target event between } t_{i-1} \text{ and } t_i \text{ sec. ago}) \\&= \int_{t_{i-1}}^{t_i} f(x)dx\end{aligned}$$

$$\sum_{i=1}^n P(\text{event starting at } t_i \text{ was targeted}) = 1$$

Algorithm2

Require: *Input: stepSize, windowSize,*

```
1: Crediter.init(windowSize)
2:  $\vec{s} \leftarrow \vec{0}$ 
3:  $\vec{f} \leftarrow \vec{0}$ 
4:  $\vec{w} \leftarrow \vec{0}$ 
5: while true do
6:   Crediter.updateTime(clockTime())
7:    $h \leftarrow \text{getHumanReinfSincePreviousTimeStep}()$ 
8:   if  $h \neq 0$  then
9:      $\vec{credFeats} \leftarrow \vec{0}$ 
10:    for all  $(\vec{f}_t, t) \in \text{Crediter.historyWindow}$  do
11:       $c_t \leftarrow \text{Crediter.assignCredit}(t)$ 
12:       $\vec{credFeats} \leftarrow \vec{credFeats} + (c_t \times \vec{f}_t)$ 
13:    end for
14:     $\text{error} \leftarrow h - (\vec{w} \cdot \vec{credFeats})$ 
15:     $\vec{w} \leftarrow \vec{w} + (\text{stepSize} \times \text{error} \times \vec{credFeats})$ 
16:  end if
17:   $\vec{s} \leftarrow \text{getStateVec}()$ 
18:   $a \leftarrow \text{argmax}_a (\vec{w} \cdot (\text{getFeatures}(\vec{s}, a)))$ 
19:   $\vec{f} \leftarrow \text{getFeatures}(\vec{s}, a)$ 
20:  takeAction(a)
21:  Crediter.updateWindow( $\vec{f}$ )
22:  wait for next time step
23: end while
```

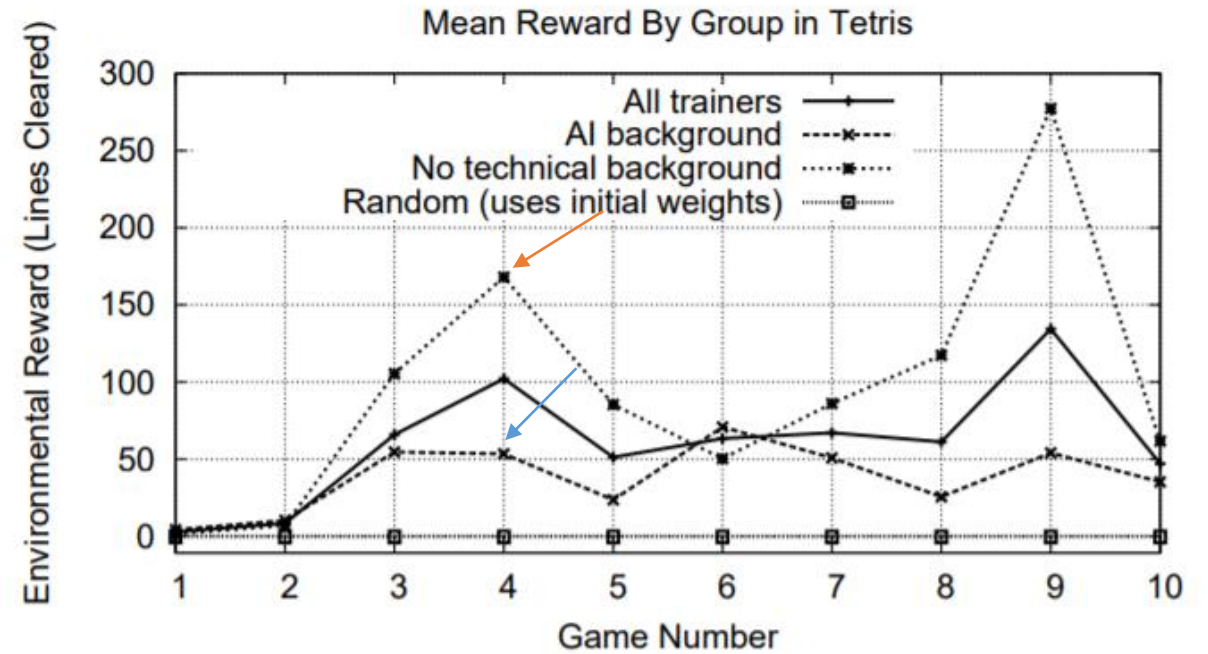
模型更新

动作选择

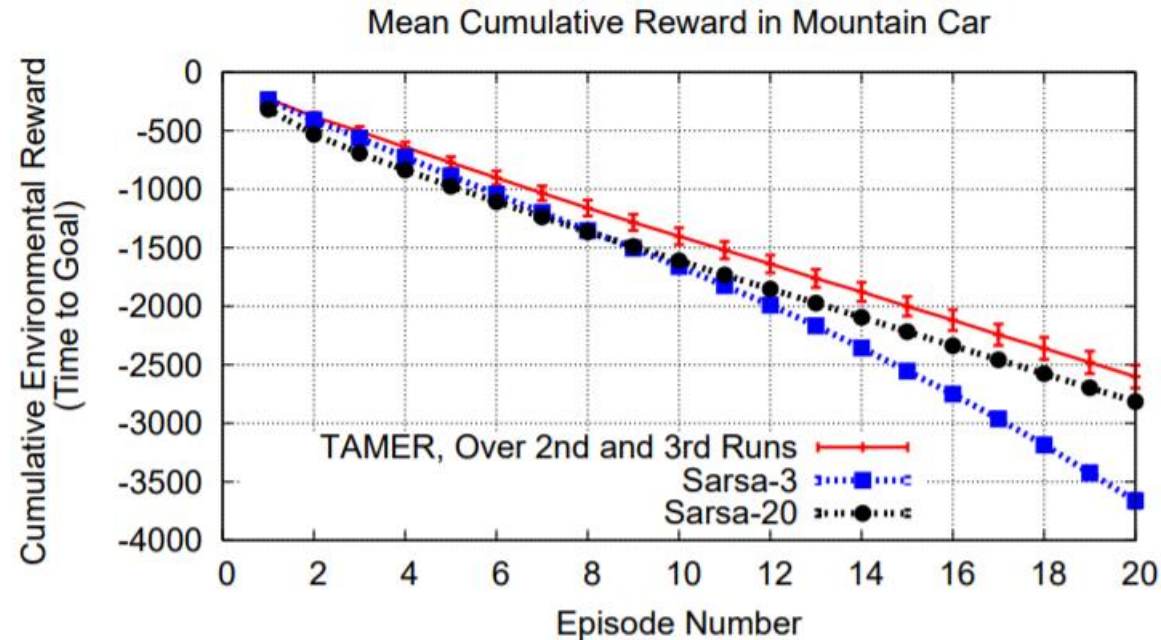
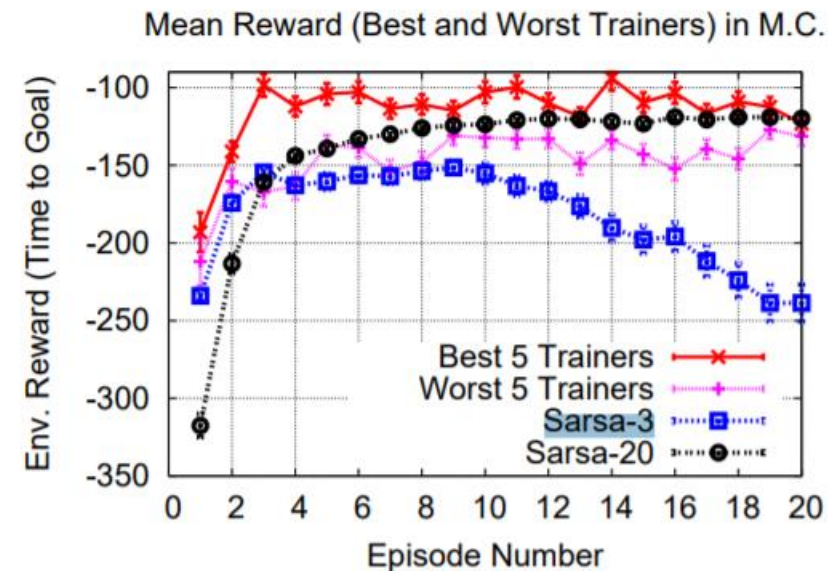
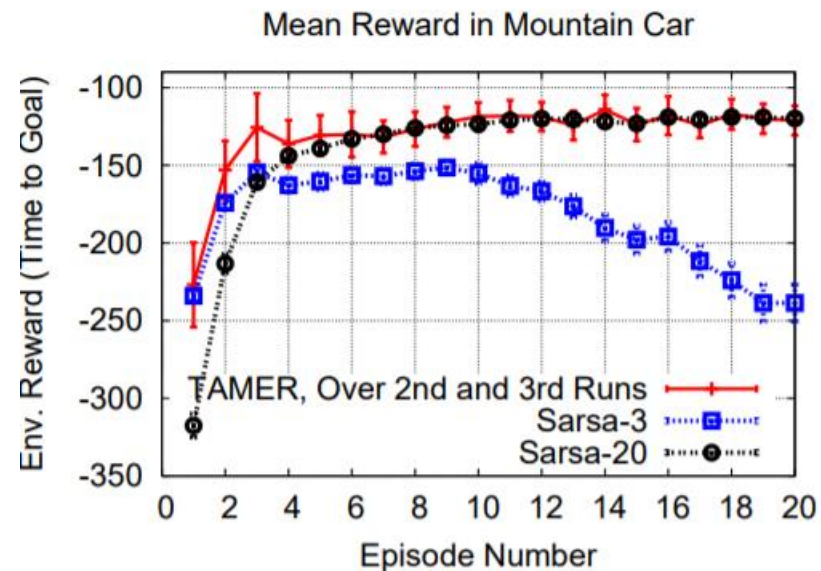
Experiment-Tetris

Table 1: Results of various Tetris agents.

Method	Mean Lines Cleared		Games for Peak
	at Game 3	at Peak	
TAMER	65.89	65.89	3
RRL-KBR [15]	5	50	120
Policy Iteration [2]	~ 0 (no learning until game 100)	3183	1500
Genetic Algorithm [5]	~ 0 (no learning until game 500)	586,103	3000
CE+RL [17]	~ 0 (no learning until game 100)	348,895	5000



Mountain Car



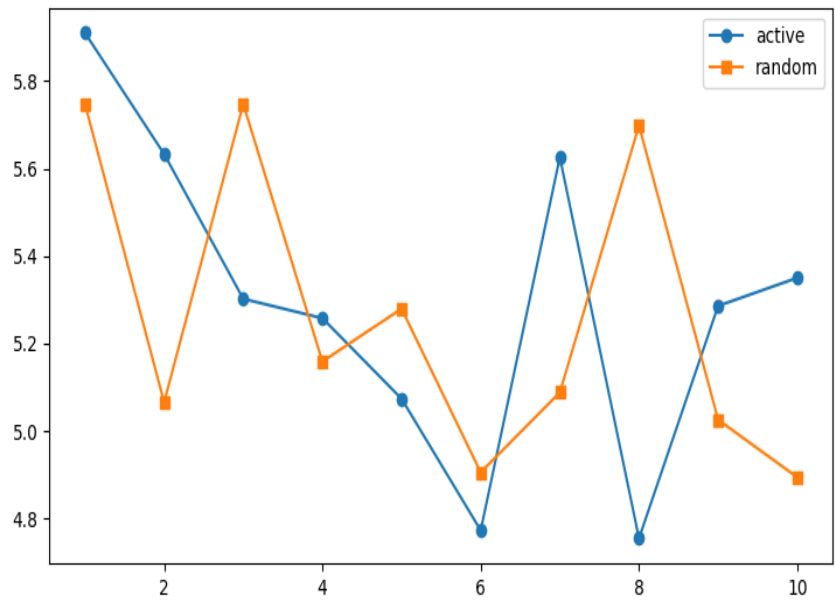
Conclusion

- works in the absence of an environmental reward function,
- reduces sample complexity
- is accessible to people who lack knowledge of computer science

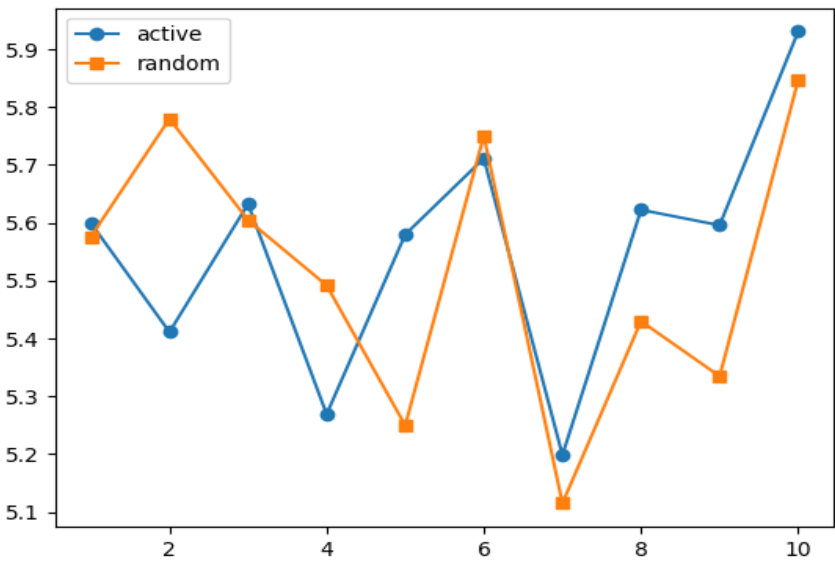
实验

逆强化学习算法：最大熵逆强化学习

X:轨迹数量
Y:reward loss

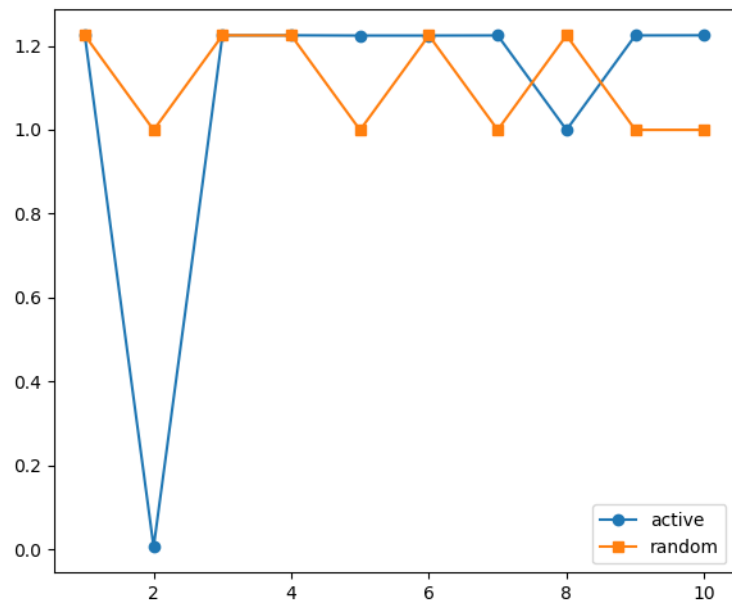


主动选初始状态以及步长vs
随机选初始状态和步长



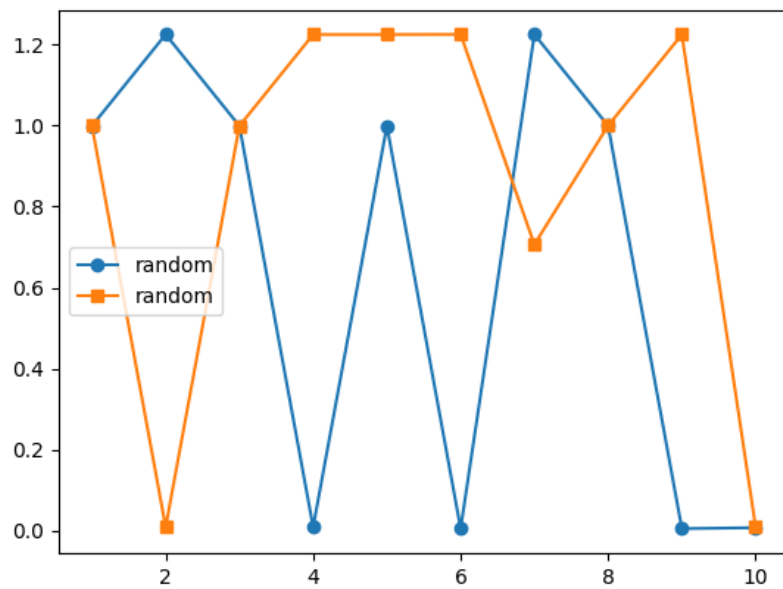
步长相同
主动选初始状态vs随机选初始状态

X:轨迹数量
Y:value loss



步长相同
主动选初始状态vs随机选初始状态

X:轨迹数量
Y:value loss



固定初始状态
vs固定选初始状态