



Artificial Intelligence and Mobility Lab



Randomized Adversarial Imitation Learning (IJCAI'19)

KICS AI Conference

Prof. Joongheon Kim

Korea University, School of Electrical Engineering
Artificial Intelligence and Mobility Laboratory

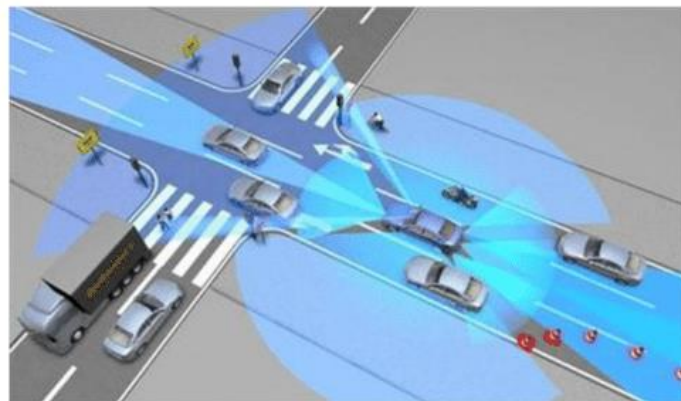
<https://joongheon.github.io>

joongheon@korea.ac.kr



- Vehicle Sensors
- Planning Algorithms
- Control Algorithms
- Deep Learning

ADAS Algorithms

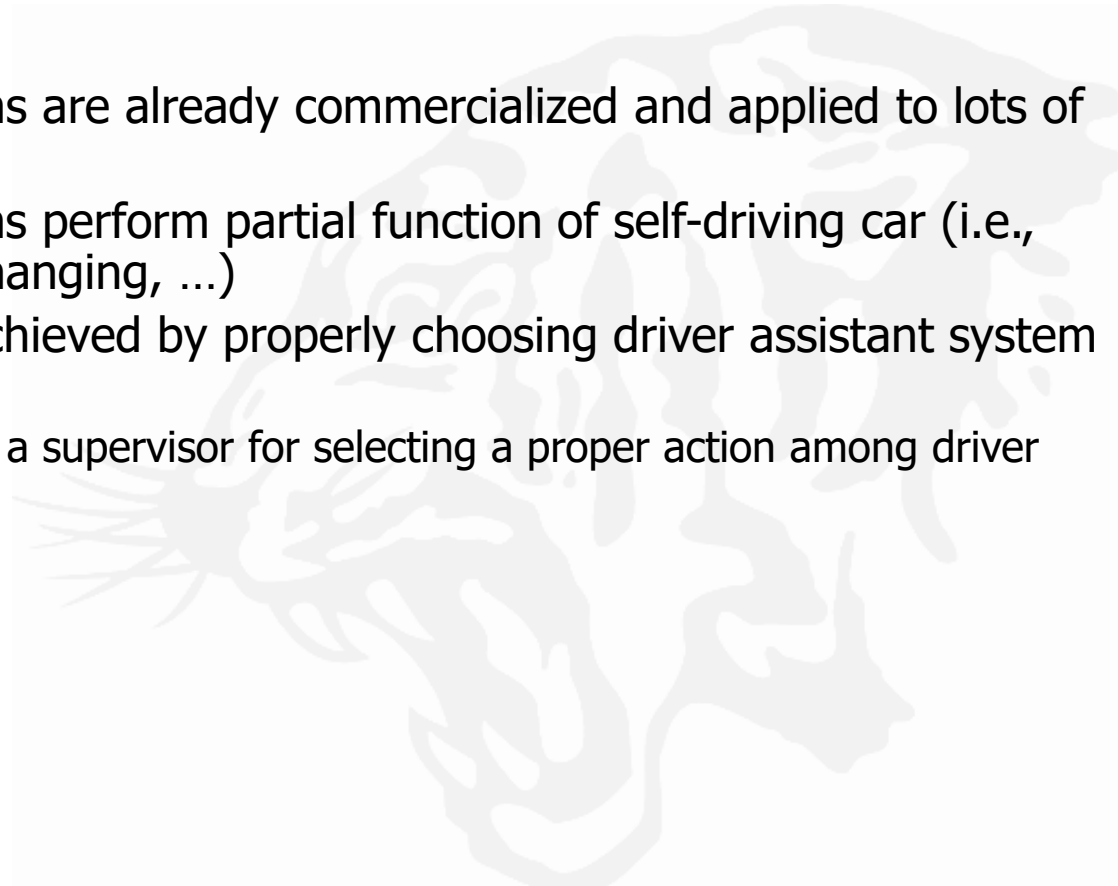


- Various ADAS
(e.g. AEB, LKAS, BSD, ESC, ...)
- Already commercialized
- Essential function for safety

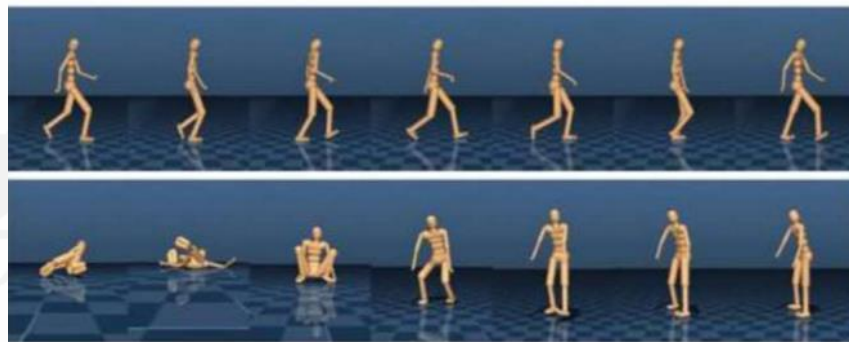
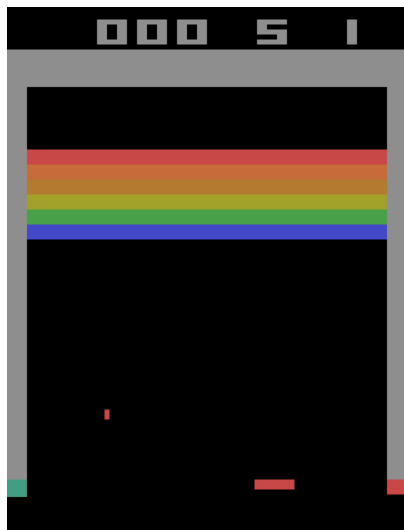
(ref) Deep Q Learning Based High Level Driving Policy Determination by Kyusik Min (IV Symposium 2018)

- Motivation

- Many driver assistance systems are already commercialized and applied to lots of vehicles.
- Many driver assistance systems perform partial function of self-driving car (i.e., lane keeping, cruising, lane changing, ...)
- Autonomous driving can be achieved by properly choosing driver assistant system in simple highway scenarios.
 - The problem is how to design a supervisor for selecting a proper action among driver assistance systems.

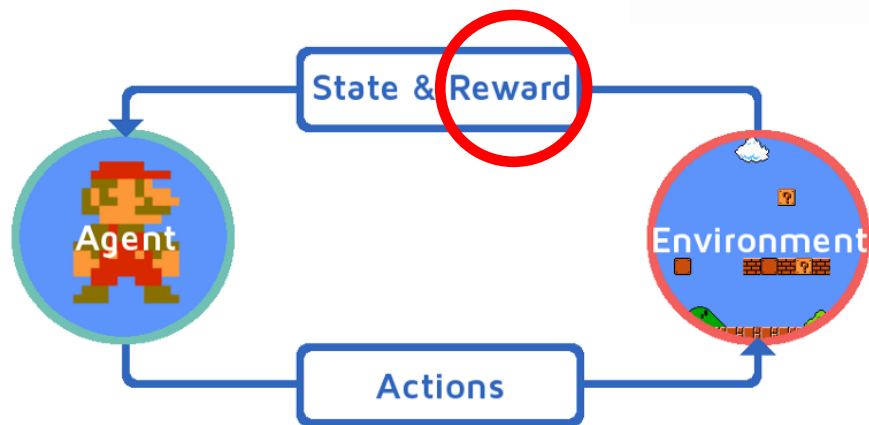


Goal : Learn policies
High-dimensional & raw observations



Super Human Performance

Challenge : Provide appropriate cost or reward signal.



- Sparse reward
- Mathematical definition

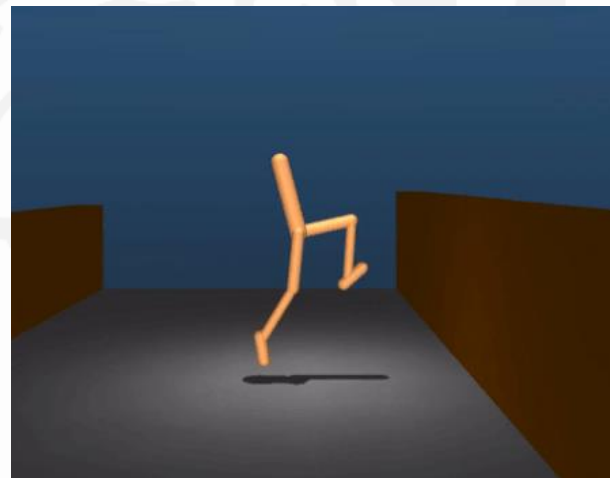
Jump is desired behavior or not?



Input : expert behavior generated by expert π_E

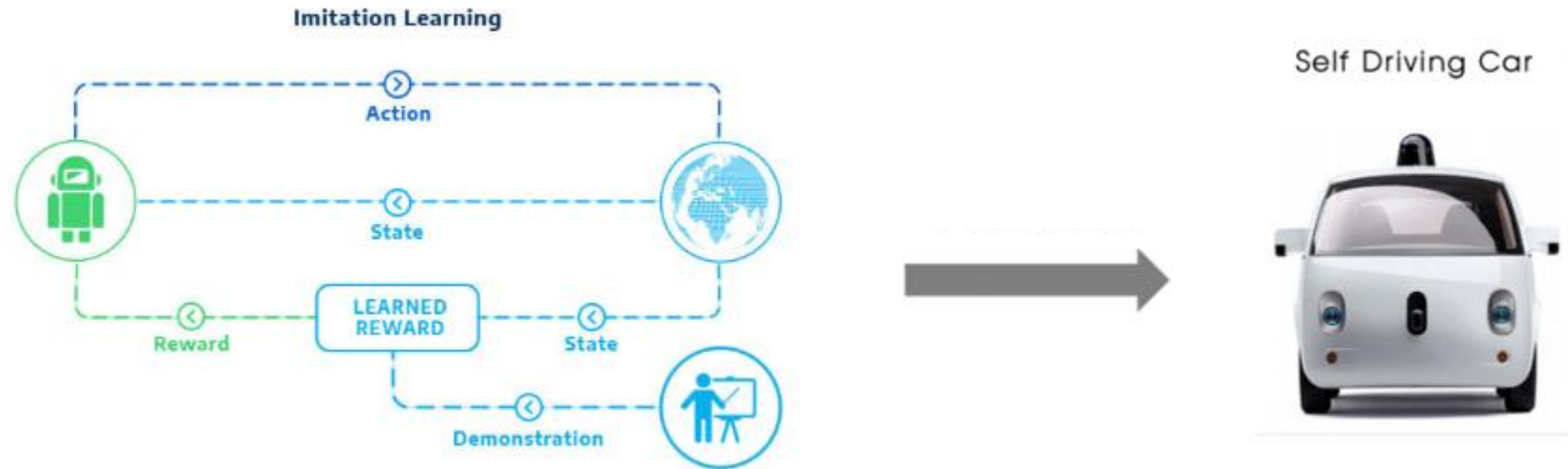
$$\left\{ \left(s_0^i, a_0^i, s_1^i, a_1^i, \dots \right) \right\}_{i=1}^N \sim \pi_E$$

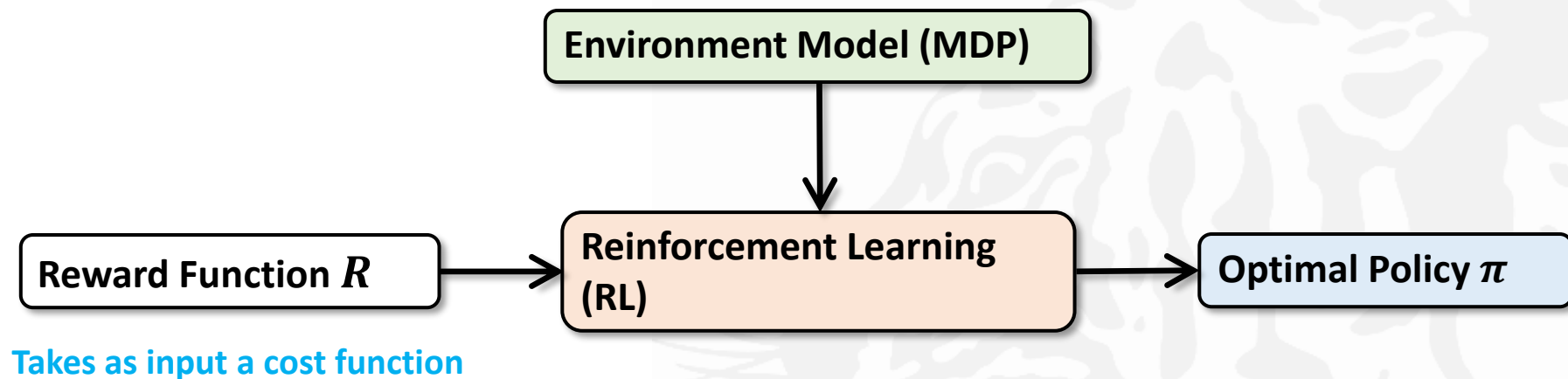
Goal : learn **cost function** or **policy**



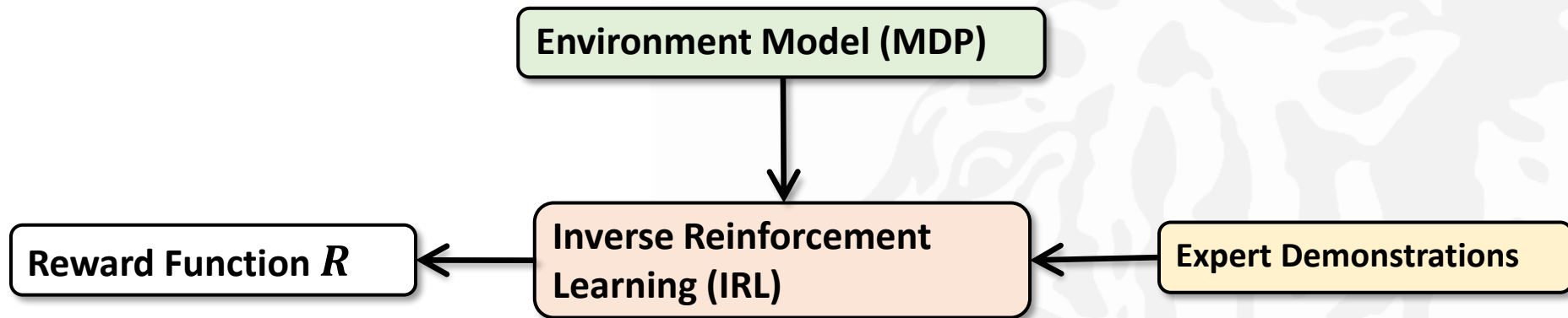
Goal

- Apply imitation learning to train supervisor of self driving car
- Enhancing safety of the self driving agent during training and testing.



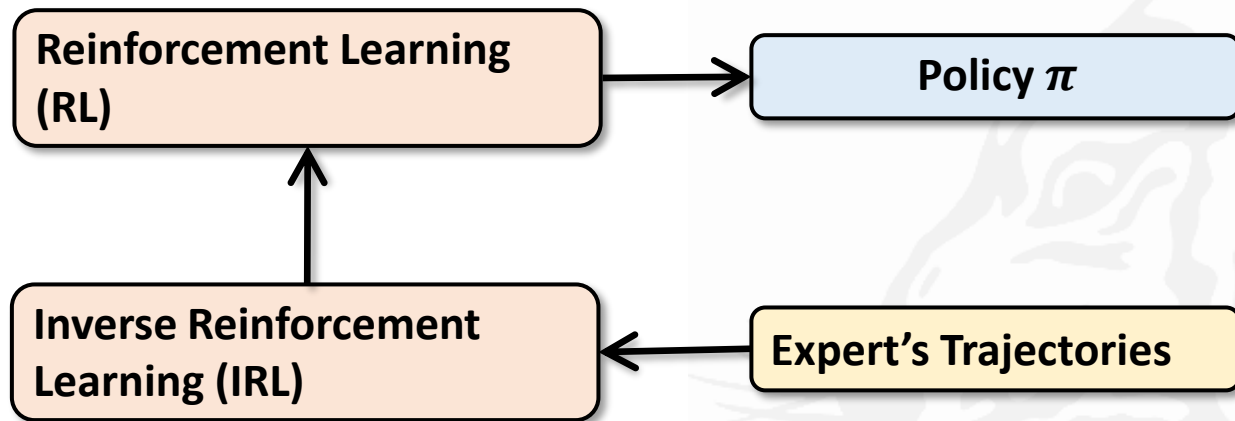


$$RL(R) = \arg \min_{\pi} \mathbb{E}_{\pi} [R(s, a)] - H(\pi)$$



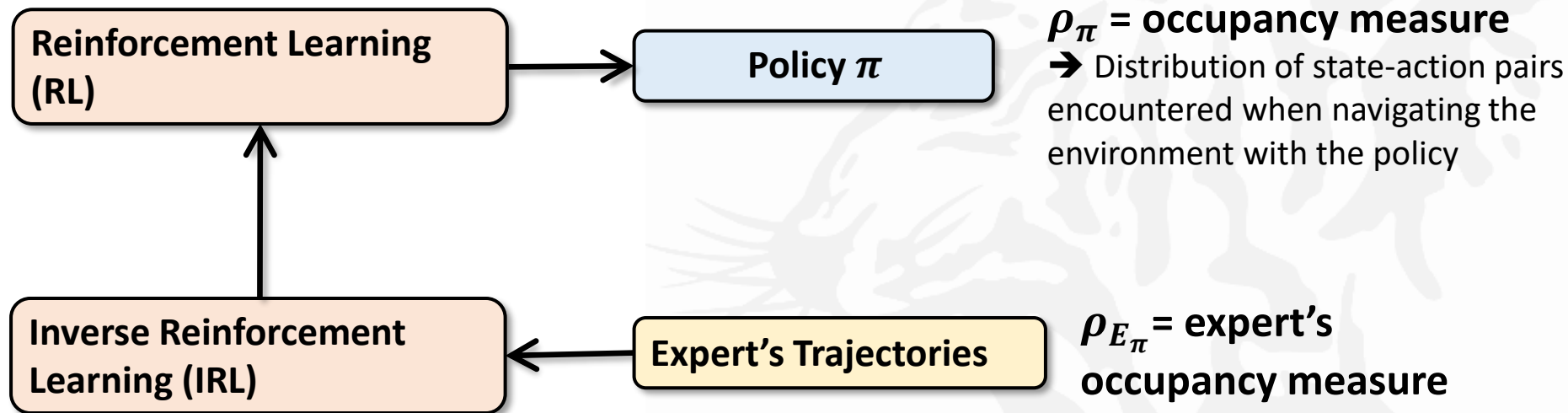
Make the expert policy uniquely optimal
with respect to that cost function.

$$\max_R \left(\min_{\pi} \mathbb{E}_{\pi} [R(s, a)] - H(\pi) \right) - \mathbb{E}_{\pi_E} [R(s, a)]$$



$$\max_R -\psi(R) + \left(\min_{\pi} \mathbb{E}_{\pi} [R(s, a)] - H(\pi) \right) - \mathbb{E}_{\pi_E} [R(s, a)]$$

Maximum Entropy Inverse Reinforcement Learning, *AAAI 2008*



[Theorem]

ψ regularized inverse reinforcement learning implicitly, seeks a policy whose occupancy measure is close to the expert's, as measured by ψ^*

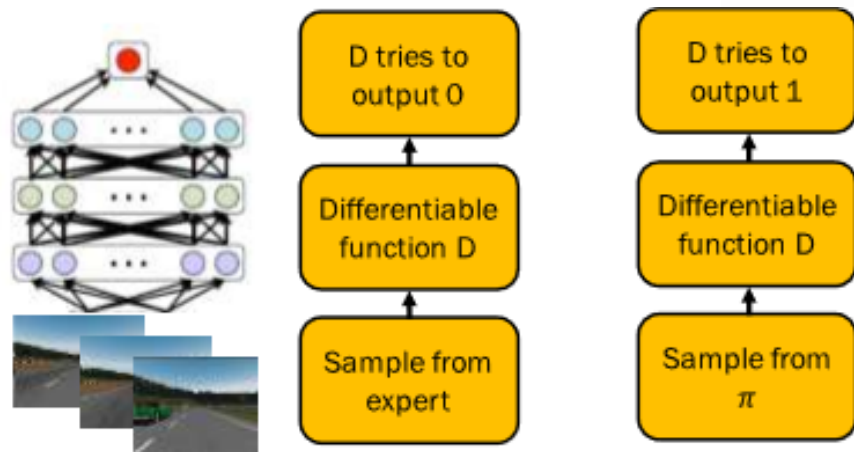
- Typical IRL finds a cost function such that the expert policy is uniquely optimal
- IRL as a procedure that tries to induce a policy that matches the expert's occupancy measure (generative model)

Generative Adversarial Imitation Learning (GAIL), *NIPS 2016*

Use this regularizer
$$\psi_{GA}(R) = \begin{cases} \mathbb{E}_{\pi_E}[g(R(s, a))] & \text{if } R < 0 \\ +\infty & \text{otherwise} \end{cases}$$

$$\min_{\pi} \psi^*(\rho_{\pi} - \rho_{E_{\pi}}) - H(\pi)$$

Generative Adversarial Imitation Learning (GAIL), *NIPS 2016*

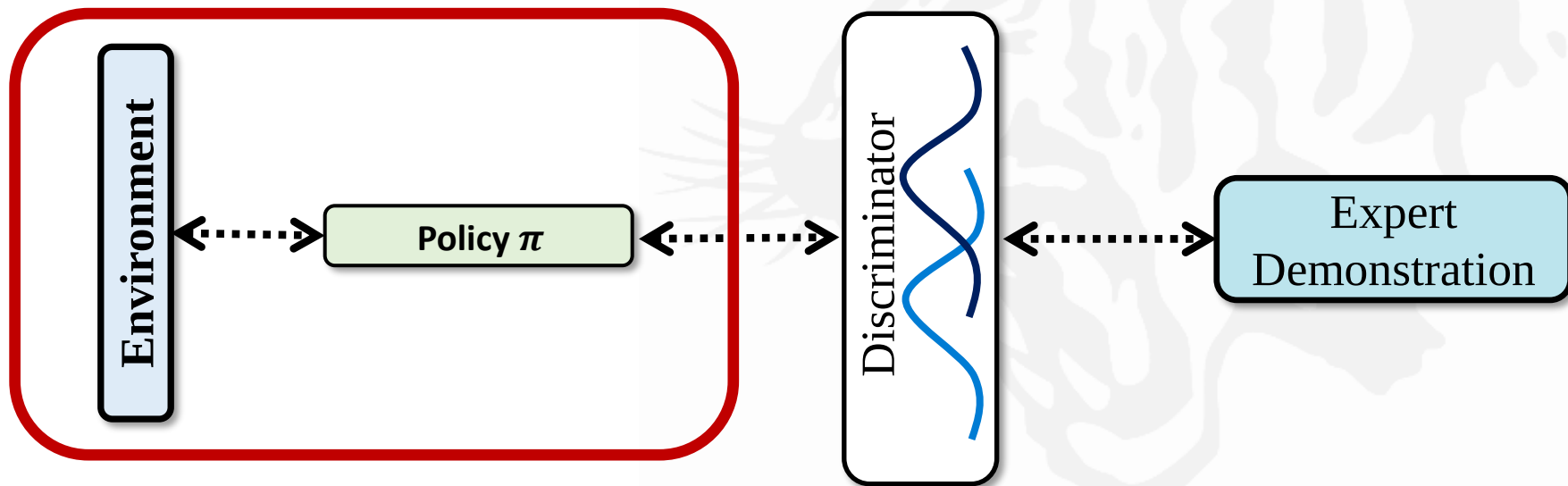


- The quality of the samples we obtain is measured by training the discriminator D.
- The policy is trained to produce behaviors that are difficult to distinguish from expert.

Adversarial Imitation Learning vis Random Search

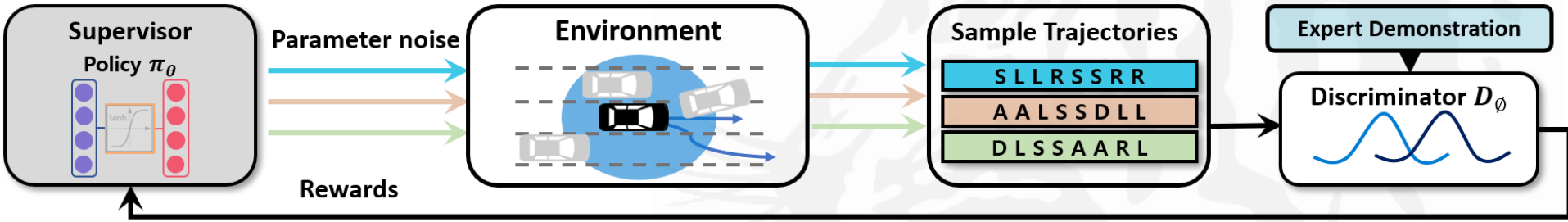
Challenge

- A lot of interaction with the environment is required to optimize the policy through GAIL framework
- Complex structure RL algorithm to train the policy.



Randomized Adversarial Imitation Learning (RAIL)

$$\text{minimize } \mathbb{E}_{\pi}[\log(D(s, a))] + \mathbb{E}_{\pi_E}[\log(1 - D(s, a))]$$



$D(s, a)$: Probability between 0 and 1

The probability that the input data sample is the expert data sample

Usually in AI:

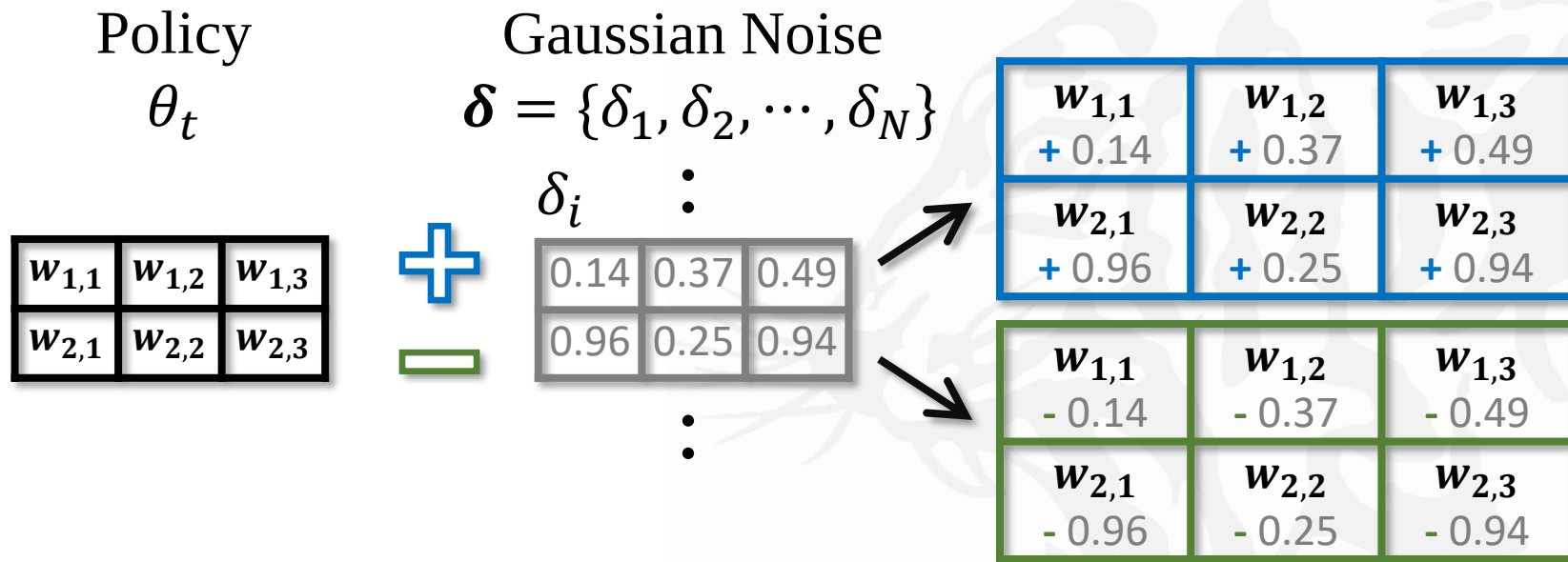
$$f'(x) = \frac{df}{dx}$$

- To update the weights of policy the gradient descent method is used.

Proposed method

$$f'(a) = \frac{f(a+h) - f(a)}{h}$$

- To update the weights of policy **the method of finite differences.**



random numbers or a matrix with random tiny values

Randomized Adversarial Imitation Learning (RAIL)

Positive perturbative weights

R_{d-pos}

$w_{1,1}$ $+d_{1,1}$	$w_{1,2}$ $+d_{1,2}$	$w_{1,3}$ $+d_{1,3}$
$w_{2,1}$ $+d_{2,1}$	$w_{2,2}$ $+d_{2,2}$	$w_{2,3}$ $+d_{2,3}$

R_{e-pos}

$w_{1,1}$ $+e_{1,1}$	$w_{1,2}$ $+e_{1,2}$	$w_{1,3}$ $+e_{1,3}$
$w_{2,1}$ $+e_{2,1}$	$w_{2,2}$ $+e_{2,2}$	$w_{2,3}$ $+e_{2,3}$

R_{f-pos}

$w_{1,1}$ $+f_{1,1}$	$w_{1,2}$ $+f_{1,2}$	$w_{1,3}$ $+f_{1,3}$
$w_{2,1}$ $+f_{2,1}$	$w_{2,2}$ $+f_{2,2}$	$w_{2,3}$ $+f_{2,3}$

R_{g-pos}

$w_{1,1}$ $+g_{1,1}$	$w_{1,2}$ $+g_{1,2}$	$w_{1,3}$ $+g_{1,3}$
$w_{2,1}$ $+g_{2,1}$	$w_{2,2}$ $+g_{2,2}$	$w_{2,3}$ $+g_{2,3}$

Negative perturbative weights

R_{d-neg}

$w_{1,1}$ $-d_{1,1}$	$w_{1,2}$ $-d_{1,2}$	$w_{1,3}$ $-d_{1,3}$
$w_{2,1}$ $-d_{2,1}$	$w_{2,2}$ $-d_{2,2}$	$w_{2,3}$ $-d_{2,3}$

R_{e-neg}

$w_{1,1}$ $-e_{1,1}$	$w_{1,2}$ $-e_{1,2}$	$w_{1,3}$ $-e_{1,3}$
$w_{2,1}$ $-e_{2,1}$	$w_{2,2}$ $-e_{2,2}$	$w_{2,3}$ $-e_{2,3}$

R_{f-neg}

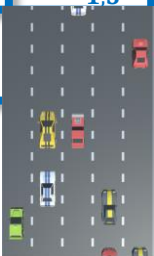
$w_{1,1}$ $-f_{1,1}$	$w_{1,2}$ $-f_{1,2}$	$w_{1,3}$ $-f_{1,3}$
$w_{2,1}$ $-f_{2,1}$	$w_{2,2}$ $-f_{2,2}$	$w_{2,3}$ $-f_{2,3}$

R_{g-neg}

$w_{1,1}$ $-g_{1,1}$	$w_{1,2}$ $-g_{1,2}$	$w_{1,3}$ $-g_{1,3}$
$w_{2,1}$ $-g_{2,1}$	$w_{2,2}$ $-g_{2,2}$	$w_{2,3}$ $-g_{2,3}$

Randomized Adversarial Imitation Learning (RAIL)

R_{d-pos}

$w_{1,1}$ $+d_{1,1}$	$w_{1,2}$ $+d_{1,2}$	$w_{1,3}$ $+d_{1,3}$
$w_{2,1}$ $+d_{2,1}$	$w_{2,2}$ $+d_{2,2}$	

R_{e-pos}

$w_{1,1}$ $+e_{1,1}$	$w_{1,2}$ $+e_{1,2}$	$w_{1,3}$ $+e_{1,3}$
$w_{2,1}$ $+e_{2,1}$	$w_{2,2}$ $+e_{2,2}$	

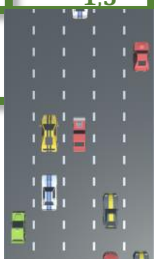
R_{f-pos}

$w_{1,1}$ $+f_{1,1}$	$w_{1,2}$ $+f_{1,2}$	$w_{1,3}$ $+f_{1,3}$
$w_{2,1}$ $+f_{2,1}$	$w_{2,2}$ $+f_{2,2}$	

R_{g-pos}

$w_{1,1}$ $+g_{1,1}$	$w_{1,2}$ $+g_{1,2}$	$w_{1,3}$ $+g_{1,3}$
$w_{2,1}$ $+g_{2,1}$	$w_{2,2}$ $+g_{2,2}$	

R_{d-neg}

$w_{1,1}$ $-d_{1,1}$	$w_{1,2}$ $-d_{1,2}$	$w_{1,3}$ $-d_{1,3}$
$w_{2,1}$ $-d_{2,1}$	$w_{2,2}$ $-d_{2,2}$	

R_{e-neg}

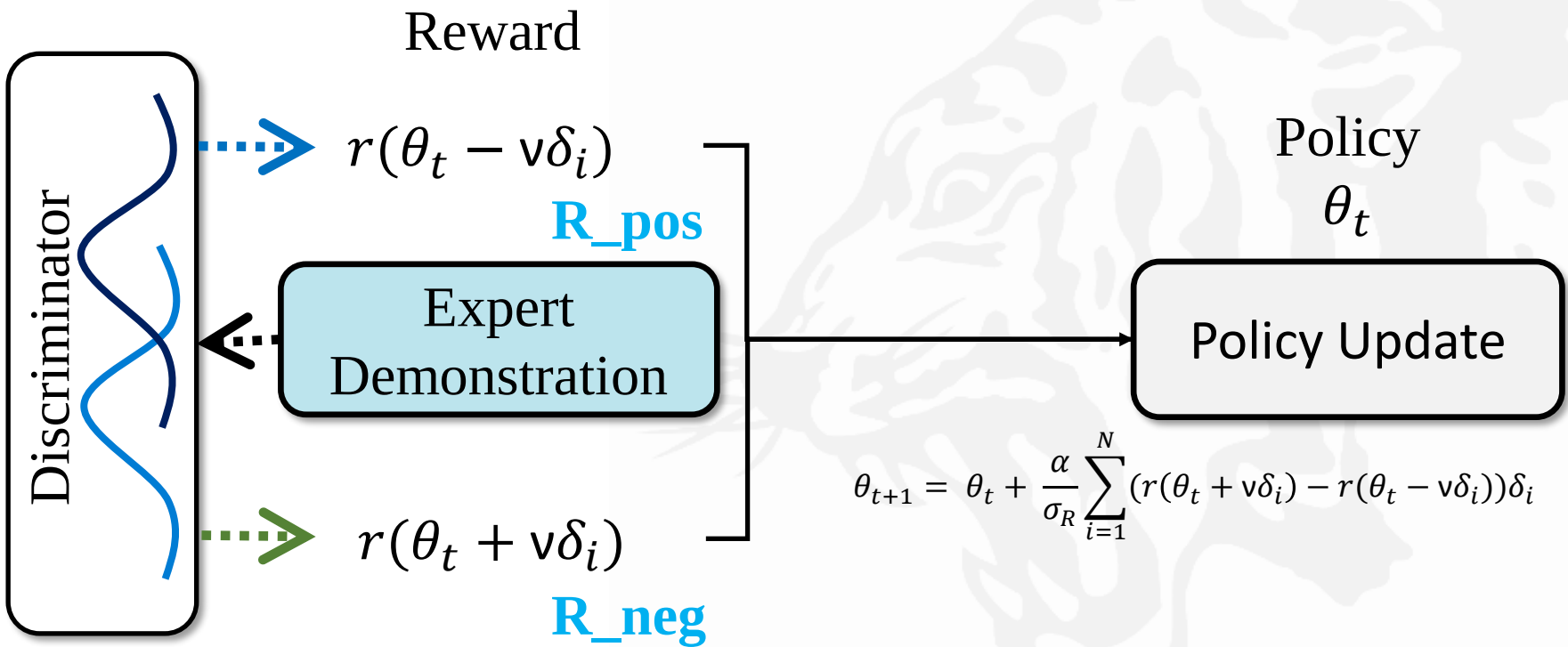
$w_{1,1}$ $-e_{1,1}$	$w_{1,2}$ $-e_{1,2}$	$w_{1,3}$ $-e_{1,3}$
$w_{2,1}$ $-e_{2,1}$	$w_{2,2}$ $-e_{2,2}$	

R_{f-neg}

$w_{1,1}$ $-f_{1,1}$	$w_{1,2}$ $-f_{1,2}$	$w_{1,3}$ $-f_{1,3}$
$w_{2,1}$ $-f_{2,1}$	$w_{2,2}$ $-f_{2,2}$	

R_{g-neg}

$w_{1,1}$ $-g_{1,1}$	$w_{1,2}$ $-g_{1,2}$	$w_{1,3}$ $-g_{1,3}$
$w_{2,1}$ $-g_{2,1}$	$w_{2,2}$ $-g_{2,2}$	



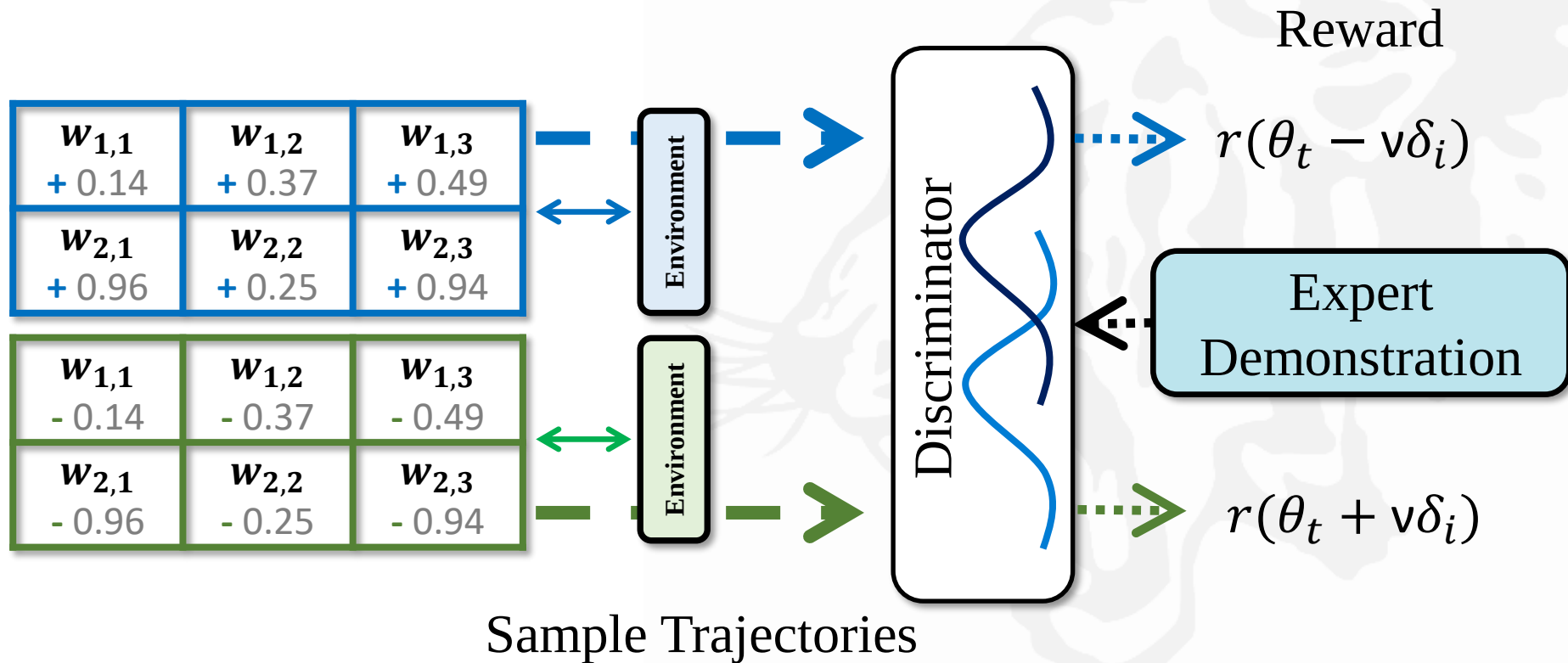
Randomized Adversarial Imitation Learning (RAIL)

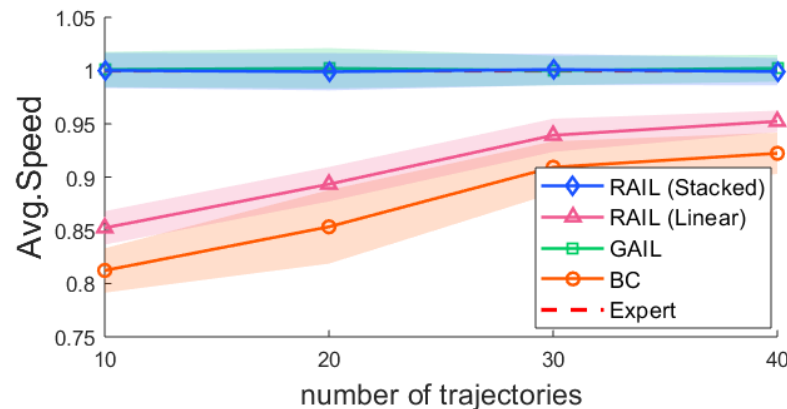
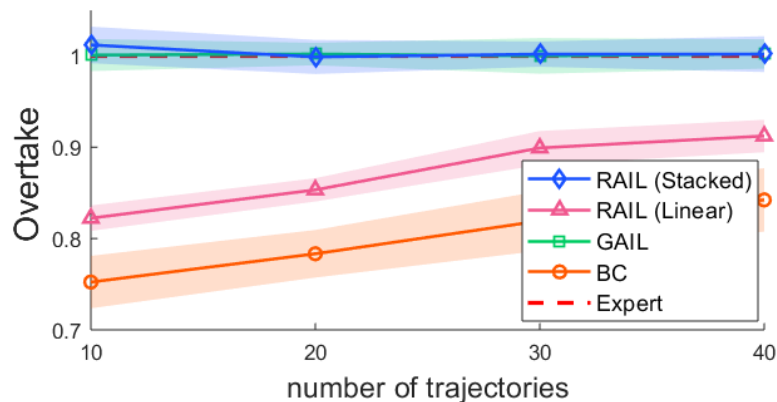
$$\begin{array}{c} \text{Policy} \\ \theta_{t+1} \end{array} = \begin{array}{c} \text{Policy} \\ \theta_t \end{array} + \left[\begin{array}{c} (R_{d\text{-}pos} - R_{d\text{-}neg}) * \begin{array}{|c|c|c|} \hline d_{1,1} & d_{1,2} & d_{1,3} \\ \hline d_{2,1} & d_{2,2} & d_{2,3} \\ \hline \end{array} + \\ (R_{e\text{-}pos} - R_{e\text{-}neg}) * \begin{array}{|c|c|c|} \hline e_{1,1} & e_{1,2} & e_{1,3} \\ \hline e_{2,1} & e_{2,2} & e_{2,3} \\ \hline \end{array} + \\ \vdots \end{array} \right]$$

Randomized Adversarial Imitation Learning (RAIL)

$$\begin{array}{c} \text{Policy} \\ \theta_{t+1} \end{array} = \begin{array}{c} \text{Policy} \\ \theta_t \end{array} + \left[\begin{array}{c} (R_{d-pos} * \begin{array}{|c|c|c|} \hline d_{1,1} & d_{1,2} & d_{1,3} \\ \hline d_{2,1} & d_{2,2} & d_{2,3} \\ \hline \end{array} - R_{d-neg} * \begin{array}{|c|c|c|} \hline d_{1,1} & d_{1,2} & d_{1,3} \\ \hline d_{2,1} & d_{2,2} & d_{2,3} \\ \hline \end{array}) + \dots \end{array} \right]$$

Randomized Adversarial Imitation Learning (RAIL)





Thank you for your attention!

- More questions?
 - joongheon@korea.ac.kr
- More details?
 - <https://joongheon.github.io/>

