

HYPERDQN: A RANDOMIZED EXPLORATION METHOD FOR DEEP REINFORCEMENT LEARNING

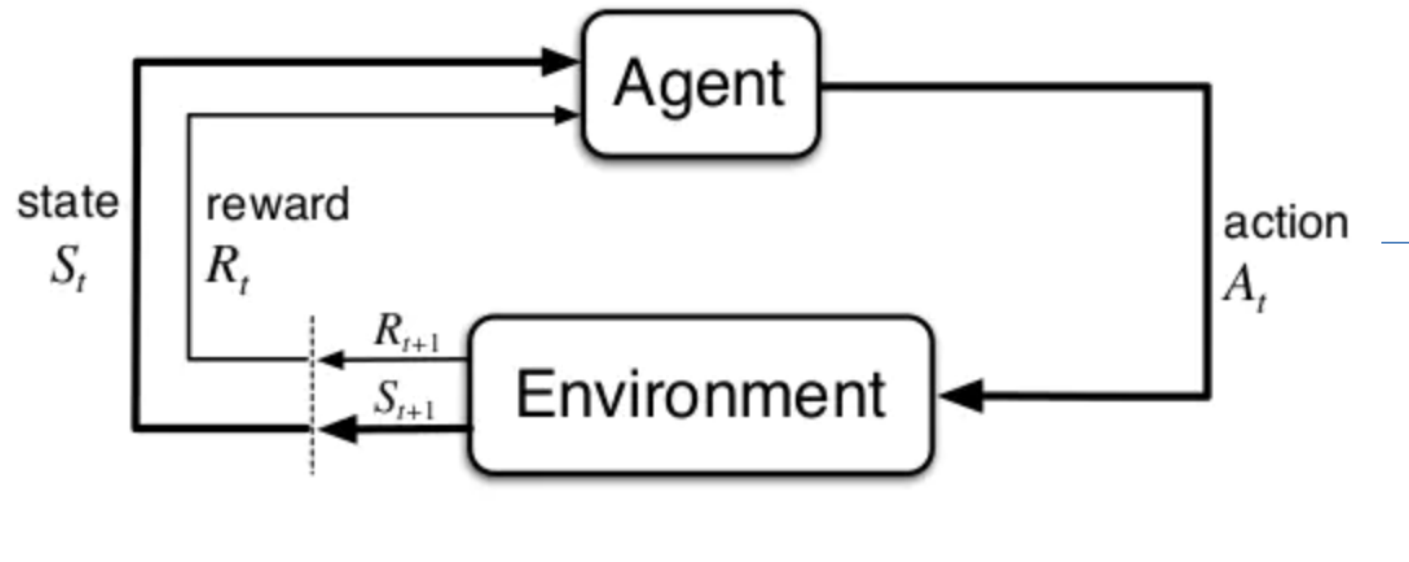
Ziniu Li¹, Yingru Li^{1†}, Yushun Zhang¹, Tong Zhang^{2†}, and Zhi-Quan Luo^{1†}

¹Shenzhen Research Institute of Big Data,
The Chinese University of Hong Kong, Shenzhen, China

²Hong Kong University of Science and Technology
{ziniuli, yingruli, yushunzhang}@link.cuhk.edu.cn,
tongzhang@ust.hk, luozq@cuhk.edu.cn

ICLR 2022

Problem: Exploration and Exploitation



- Exploitation → higher immediate reward & local optimum
- Exploration → sacrifice immediate reward & wider space

1. ϵ -greedy : Select greedy action with probability of $(1-\epsilon)$
Select action randomly with probability of ϵ
 ϵ decreases with the increase of the agent step

2. Boltzmann distribution : Select action according to the Boltzmann distribution of Q-value (softmax)

3. Upper confidence bounds(UCB) : $A_t^{UCB} \doteq \arg \max_a \left[Q_t + c \sqrt{\frac{\ln t}{N_t(a)}} \right]$

$N_t(a)$: The number of action a selected before time t

4. Thomson Sample : posterior sampling

- (1) Set up Q-value distribution estimates for each action
- (2) Sample randomly from each distribution to get Q-value, and select the action with the maximum Q-value
- (3) Update distribution parameters based on reward

- The upper regret bound is reduced

$$\text{Regret}(T, \mathcal{M}) = \sum_{l=0}^{T/H-1} \mathbb{E}_{\mathcal{M}} \left[V_0^*(s_{l0}) - \sum_{h=0}^H r_{lh} \right]$$

- The randomness in posterior sampling could yield positive bias, which boosts optimistic behaviors
- the property of temporal extended exploration(deep exploration)

Algorithm 2 RLSVI with greedy action

Input: Features $\Phi_0, \dots, \Phi_{H-1}$; $\sigma > 0, \lambda > 0$

- 1: **for** $l=0, 1, \dots$ **do**
 - 2: Compute $\tilde{\theta}_{l0}, \dots, \tilde{\theta}_{l, H-1}$ using Algorithm 1
 - 3: Observe s_{l0}
 - 4: **for** $h=0, \dots, H-1$ **do**
 - 5: Sample $a_{lh} \in \arg\max_{\alpha \in \mathcal{A}} \left(\Phi_h \tilde{\theta}_{lh} \right) (s_{lh}, \alpha)$
 - 6: Observe r_{lh} and $s_{l, h+1}$
 - 7: **end for**
 - 8: Observe r_{lH}
 - 9: **end for**
-

Step1: Given initial value of θ ($l=0$)

Step2: Q-value comes from $\Phi\theta$, and select action with greedy strategy

Step3: Update θ according to Bayesian regression

Step4: Repeat steps 2 and 3 for each episode

Algorithm 1 Randomized Least-Squares Value Iteration

Input: Data $\Phi_0(s_{i0}, a_{i0}), r_{i0}, \dots, \Phi_{H-1}(s_{iH-1}, a_{iH-1}), r_{iH}$: $i < L$, Parameters $\lambda > 0, \sigma > 0$

Output: $\tilde{\theta}_{l0}, \dots, \tilde{\theta}_{l, H-1}$

- 1: **for** $h = H-1, \dots, 1, 0$ **do**
- 2: Generate regression problem $A \in \mathbb{R}^{l \times K}, b \in \mathbb{R}^l$:

$$A \leftarrow \begin{bmatrix} \Phi_h(s_{0h}, a_{0h}) \\ \vdots \\ \Phi_h(s_{l-1, h}, a_{l-1, h}) \end{bmatrix}$$

$$b_i \leftarrow \begin{cases} r_{ih} + \max_{\alpha} \left(\Phi_{h+1} \tilde{\theta}_{l, h+1} \right) (s_{i, h+1}, \alpha) & \text{if } h < H-1 \\ r_{ih} + r_{i, h+1} & \text{if } h = H-1 \end{cases}$$

- 3: Bayesian linear regression for the value function

$$\bar{\theta}_{lh} \leftarrow \frac{1}{\sigma^2} \left(\frac{1}{\sigma^2} A^\top A + \lambda I \right)^{-1} A^\top b$$

$$\Sigma_{lh} \leftarrow \left(\frac{1}{\sigma^2} A^\top A + \lambda I \right)^{-1}$$

- 4: Sample $\tilde{\theta}_{lh} \sim N(\bar{\theta}_{lh}, \Sigma_{lh})$ from Gaussian posterior
 - 5: **end for**
-

Bayesian Linear Regression

conditions: $y = \langle \theta^*, x \rangle + \omega^*$, $\omega^* \sim N(0, \sigma_\omega^2)$ prior distribution $p_0(\theta^*) \sim N(\bar{\theta}_p, \sigma_p^2)$ Dataset $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$

According to Bayes formula:

$$p(\theta^*|\mathcal{D}) = \frac{p(\mathcal{D}|\theta^*)p_0(\theta^*)}{p(\mathcal{D})} \Rightarrow p(\theta^*|\mathcal{D}) \propto p(\mathcal{D}|\theta^*)p_0(\theta^*) \Rightarrow p(\mathcal{D}|\theta^*) = p(Y|\theta^*, X) \sim N(X\theta^*, \sigma_\omega^2)$$

Conjugate prior: For some likelihood functions, its prior and posterior have the same distribution

Multi-dimensional gaussian distribution:

$$f(x) = \frac{1}{(2\pi)^{\frac{n}{2}} |\Sigma|^{\frac{1}{2}}} \exp \left\{ -\frac{1}{2} (X - \mu)^T \Sigma^{-1} (X - \mu) \right\}, \quad \Sigma: \text{Covariance matrix}$$

So we can assume that the distribution: $p(\theta^*|\mathcal{D}) \sim N(\mu, \Sigma)$

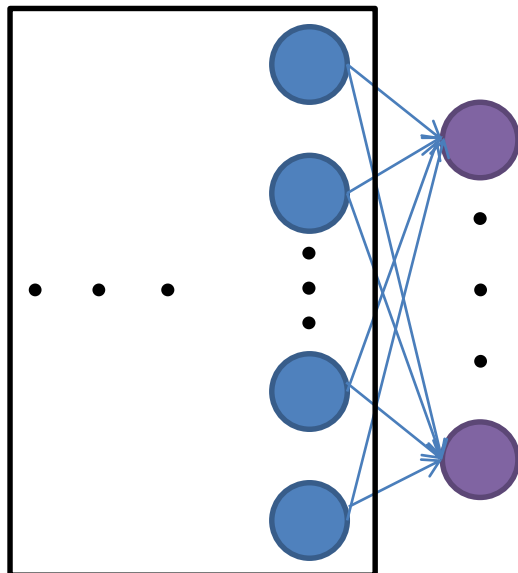
For an ordinary multidimensional Gaussian distribution:

$$-\frac{1}{2} (X - \mu)^T \Sigma^{-1} (X - \mu) = -\frac{1}{2} (\mathbf{X}^T \Sigma^{-1} \mathbf{X} - \mathbf{X}^T \Sigma^{-1} \mu - \mu^T \Sigma^{-1} \mathbf{X} + \mu^T \Sigma^{-1} \mu)$$

$$\begin{aligned} p(\mathcal{D}|\theta^*) &\propto p(\mathcal{D}|\theta^*)p_0(\theta^*) \propto \exp \left\{ -\frac{1}{2\sigma_\omega^2} (Y^T - \theta^{*T} X^T) (Y - X\theta^*) \right\} \exp \left\{ -\frac{1}{2\sigma_p^2} (\theta^* - \bar{\theta}_p)^T (\theta^* - \bar{\theta}_p) \right\} \\ &= \exp \left\{ -\frac{1}{2\sigma_\omega^2} (Y^T Y - \theta^{*T} X^T Y - Y^T X \theta^* + \theta^{*T} X^T X \theta^*) - \frac{1}{2\sigma_p^2} (\theta^* - \bar{\theta}_p)^T (\theta^* - \bar{\theta}_p) \right\} \\ &= \exp \left\{ -\frac{1}{2\sigma_\omega^2} (Y^T Y - \theta^{*T} X^T Y - Y^T X \theta^* + \theta^{*T} X^T X \theta^*) - \frac{1}{2} (\theta^* - \bar{\theta}_p)^T (\theta^* - \bar{\theta}_p) \right\} \\ &= \exp \left\{ -\frac{1}{2} [\theta^{*T} (\sigma_\omega^{-2} X^T X + \sigma_p^{-2} I) \theta^*] + \frac{1}{2} [\theta^{*T} (\sigma_\omega^{-2} X^T Y + \sigma_p^{-2} \bar{\theta}_p) + (Y^T X \sigma_\omega^{-2} + \bar{\theta}_p^T \sigma_p^{-2}) \theta^*] + c \right\} \end{aligned}$$

Therefore $\Sigma = \sigma_\omega^{-2} X^T X + \sigma_p^{-2} I$ $\mu = (\sigma_\omega^{-2} X^T X + \sigma_p^{-2} I)^{-1} (\sigma_\omega^{-2} X^T Y + \sigma_p^{-2} \bar{\theta}_p)$

➤ Provided feature Φ



Features Φ Q

➤ The computational complexity

fixed ϕ : $\Phi_K = \Phi_{K-1} + \phi(x_K)\phi(x_K)^\top$ with $\Phi_0 = \mathbf{0}$,

changing ϕ_K : $\Phi_K := \sum_{\ell=1}^K \phi_K(x_\ell)\phi_K(x_\ell)^\top$, $\Phi_{K-1} := \sum_{\ell=1}^{K-1} \phi_{K-1}(x_\ell)\phi_{K-1}(x_\ell)^\top, \dots$

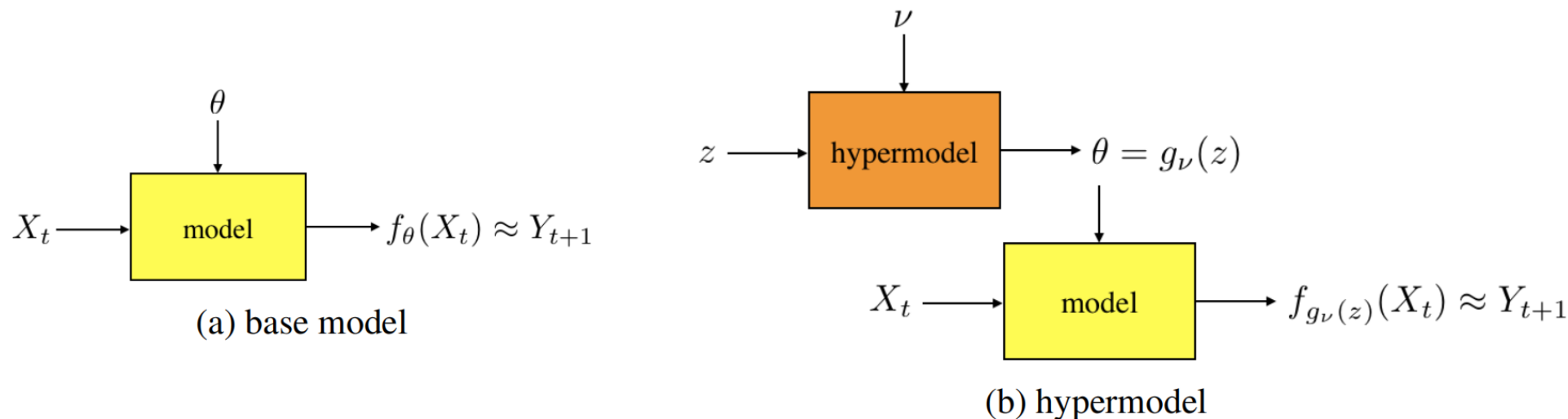
➡ Covariance matrix

Algorithm 2 RLSVI with greedy action

Input: Features $\Phi_0, \dots, \Phi_{H-1}$; $\sigma > 0$, $\lambda > 0$

```

1: for  $l=0,1,\dots$  do
2:   Compute  $\tilde{\theta}_{l0}, \dots, \tilde{\theta}_{l,H-1}$  using Algorithm 1
3:   Observe  $s_{l0}$ 
4:   for  $h=0, \dots, H-1$  do
5:     Sample  $a_{lh} \in \arg\max_{\alpha \in \mathcal{A}} (\Phi_h \tilde{\theta}_{lh})(s_{lh}, \alpha)$ 
6:     Observe  $r_{lh}$  and  $s_{l,h+1}$ 
7:   end for
8:   Observe  $r_{lH}$ 
9: end for
    
```

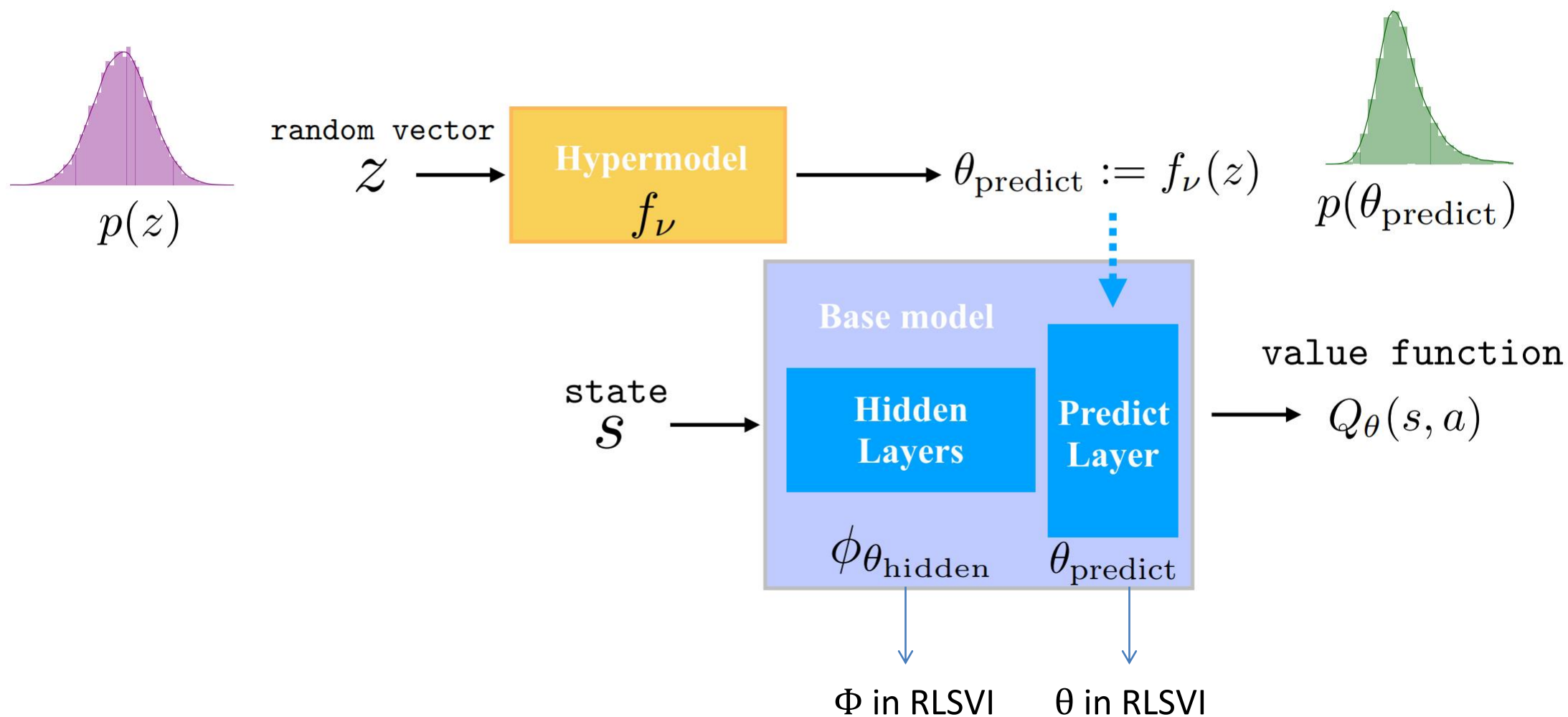


Theorem 1 Let p_z be the unit Gaussian distribution in $\mathbb{R}^K$. For all $\epsilon > 0$, $\delta > 0$, $B > 0$, and probability measures μ over $L_\infty(\mathcal{X}, B)$, there exists a transport map H from p_z to μ , a neural network $f_\theta : \mathcal{X} \mapsto \mathbb{R}$ with a linear output node and ReLU hidden nodes, and a linear hypermodel $g_\nu : \mathcal{Z} \mapsto \mathbb{R}^{N_\theta}$ such that

$$\|f_{g_\nu(z)} - f^*\|_\infty \leq \epsilon,$$

with probability at least $1 - \delta$, for some $\nu \in \mathbb{R}^{N_\nu}$, where $f^* = H(z)$.

A linear hypermodel can represent essentially any distribution over functions



Loss function

$$L(\nu; \mathcal{D}) = \int_z p(z) \left[\sum_{(x,y,\xi) \in \mathcal{D}} \left(y + \underbrace{\sigma_\omega z^\top \xi}_{(a)} - \underbrace{(g_{f_{\nu_{\text{prior}}}(z)}(x) + g_{f_\nu(z)}(x))}_{(b)} \right)^2 + \underbrace{\frac{\sigma_\omega^2}{\sigma_p^2} \|f_\nu(z)\|^2}_{(c)} \right] (dz), \quad (4.1)$$

$$\min_{\nu, \theta_{\text{hidden}}} \int_z p(z) \left[\sum_{(s,a,r,\xi,s') \in \mathcal{D}} \left(Q_{\text{target}}(s', z) + \sigma_\omega z^\top \xi - Q_{\text{prediction}}(s, a, z) \right)^2 + \frac{\sigma_\omega^2}{\sigma_p^2} \|f_\nu(z)\|^2 \right] (dz), \quad (4.2)$$

where

$$Q_{\text{prediction}}(s, a, z) = Q_{\theta_{\text{prior}}, f_{\nu_{\text{prior}}}(z)}(s, a) + Q_{\theta_{\text{hidden}}, f_\nu(z)}(s, a), \quad (4.3)$$

$$Q_{\text{target}}(s', z) = r + \gamma \max_{a'} \left[Q_{\theta_{\text{prior}}, f_{\nu_{\text{prior}}}(z)}(s', a') + Q_{\bar{\theta}_{\text{hidden}}, f_{\bar{\nu}}(z)}(s', a') \right].$$

$p(z)$: Gaussian distribution

ξ : a random vector independently sampled from the unit hypersphere

$\sigma_\omega z^\top \xi$: is an artificial noise term exerted on the label y

Generate $\mathbf{x} \sim N(0, I)$

Normalize $\mathbf{x}$

$\frac{\sigma_\omega^2}{\sigma_p^2} \|f_\nu(z)\|^2$: a regularization term

$$L(\nu; \mathcal{D}) = \int_{\mathcal{Z}} p(z) \left[\sum_{(x, y, \xi) \in \mathcal{D}} \left(y + \underbrace{\sigma_{\omega} z^{\top} \xi}_{(a)} - \underbrace{(g_{f_{\nu_{\text{prior}}}(z)}(x) + g_{f_{\nu}(z)}(x))}_{(b)} \right)^2 + \underbrace{\frac{\sigma_{\omega}^2}{\sigma_p^2} \|f_{\nu}(z)\|^2}_{(c)} \right] (dz), \quad (4.1)$$

Assumption 1. Suppose the data generation follows $y = x^{\top} \theta^* + \omega^*$, $\omega^* \sim \mathcal{N}(0, \sigma_{\omega}^2)$ and the prior distribution over θ^* is $\mathcal{N}(\bar{\theta}_p, \sigma_p^2 I)$. Furthermore, assume the base model is linear, i.e., $g_{f_{\nu}(z)}(x) = x^{\top} f_{\nu}(z)$ and $g_{f_{\nu_{\text{prior}}}(z)}(x) = x^{\top} f_{\nu_{\text{prior}}}(z)$. Moreover, assume the hypermodel is also linear, i.e., $f_{\nu}(z) = \nu_w^{\top} z + \nu_b$ and $f_{\nu_{\text{prior}}}(z) = \nu_w^{\text{prior}\top} z + \nu_b^{\text{prior}}$; $\nu = (\nu_w, \nu_b)$ and $\nu_{\text{prior}} = (\nu_w^{\text{prior}}, \nu_b^{\text{prior}})$.

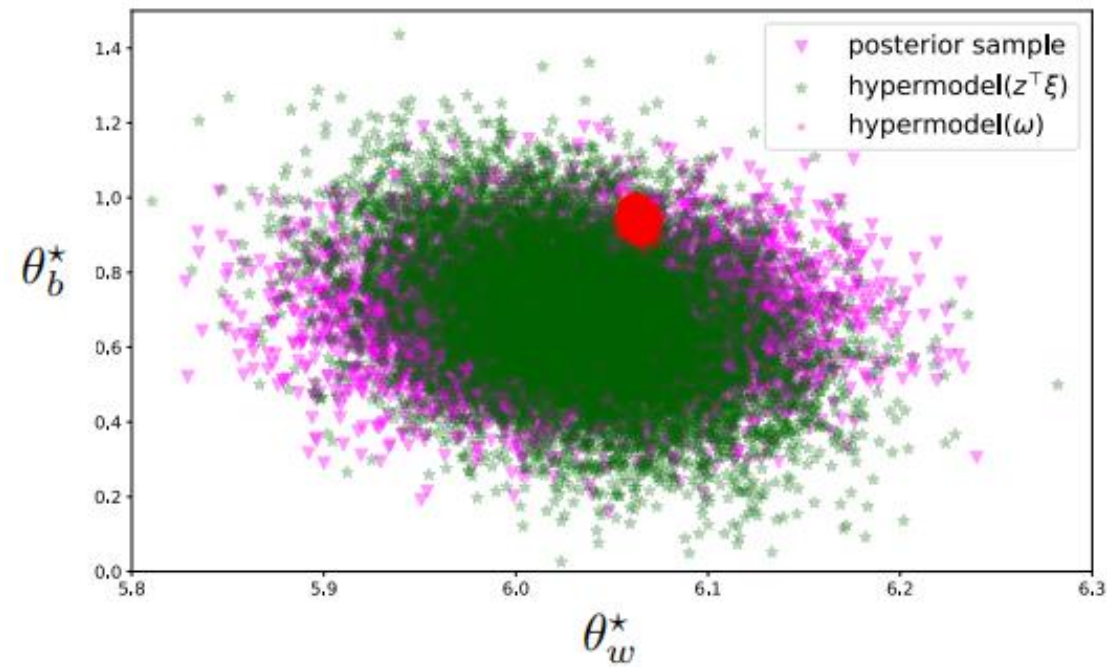
Theorem 1 (Formal statement). Under Assumption 1, set $\nu_w^{\text{prior}} = \sigma_p I$ and $\nu_b^{\text{prior}} = \bar{\theta}_p$. Let $\nu^* = (\nu_w^*, \nu_b^*)$ be the optimal solution of (4.1) conditioned on specific realizations of ξ , then $\theta := f_{\nu_{\text{prior}}}(z) + f_{\nu^*}(z) \sim \mathcal{N}(\nu_b^{\text{prior}} + \nu_b^*, (\nu_w^{\text{prior}} + \nu_w^*)^{\top} (\nu_w^{\text{prior}} + \nu_w^*))$ with

$$\nu_b^{\text{prior}} + \nu_b^* = \mathbb{E}[\theta^* | \mathcal{D}], \quad (\nu_w^{\text{prior}} + \nu_w^*)^{\top} (\nu_w^{\text{prior}} + \nu_w^*) = \text{Cov}[\theta^* | \mathcal{D}] + \text{err}(\Xi),$$

where

$$\text{err}(\Xi) := \text{Cov}[\theta^* | \mathcal{D}] \left(\frac{1}{\sigma_{\omega}^2} \sum_{(x, \xi) \neq (x', \xi') \in \mathcal{D}} x \xi^{\top} \xi' x'^{\top} + \frac{1}{\sigma_{\omega} \sigma_p} \sum_{(x, \xi) \in \mathcal{D}} (x \xi^{\top} + \xi x^{\top}) \right) \text{Cov}[\theta^* | \mathcal{D}]$$

Furthermore, the error term satisfies $\mathbb{E}_{\Xi}[\text{err}(\Xi)] = 0$.



z-dependent noise is indispensable

Algorithm 2 HyperDQN

```

1: agent step  $n \leftarrow 0$ , train step  $\ell \leftarrow 0$ .
2: for episode  $k = 0, 1, 2, \dots$  do
3:   generate a random vector  $z \sim \mathcal{N}(0, I)$ . ▷ Sampling
4:   instantiate the  $Q$ -value function  $Q_\theta(s, a, z)$ .
5:   for stage  $t = 0, 1, 2, \dots, T - 1$  do ▷ Interaction
6:     observe state  $s_t$ .
7:     take the greedy action  $a \leftarrow \operatorname{argmax}_a Q_\theta(s_t, a, z)$ .
8:     receive the next state  $s_{t+1}$  and reward  $r(s_t, a_t)$ .
9:     sample  $\xi$  uniformly from unit hypersphere.
10:    store  $(s_t, a_t, r, \xi, s_{t+1})$  into the replay buffer  $\mathcal{D}$ .
11:    agent step  $n \leftarrow n + 1$ .
12:    if  $\text{mod}(\text{agent step } n, \text{train frequency } M) == 0$  then ▷ Update
13:      sample a mini-batch  $\tilde{\mathcal{D}}$  of  $(s, a, r, \xi, s')$  from the replay buffer  $\mathcal{D}$ .
14:      sample  $N$  random vectors  $z'_i$ :  $\tilde{\mathcal{Z}} = \{z'_i\}_{i=1}^N$ .
15:      optimize  $\nu$  and  $\theta_{\text{hidden}}$  using the empirical loss function (A.6) with  $\tilde{\mathcal{D}}$  and  $\tilde{\mathcal{Z}}$ .
16:      train step  $\ell \leftarrow \ell + 1$ .
17:    end if
18:    if  $\text{mod}(\text{train step } \ell, \text{target update frequency } G) == 0$  then
19:      update the target network.
20:    end if
21:  end for
22: end for

```

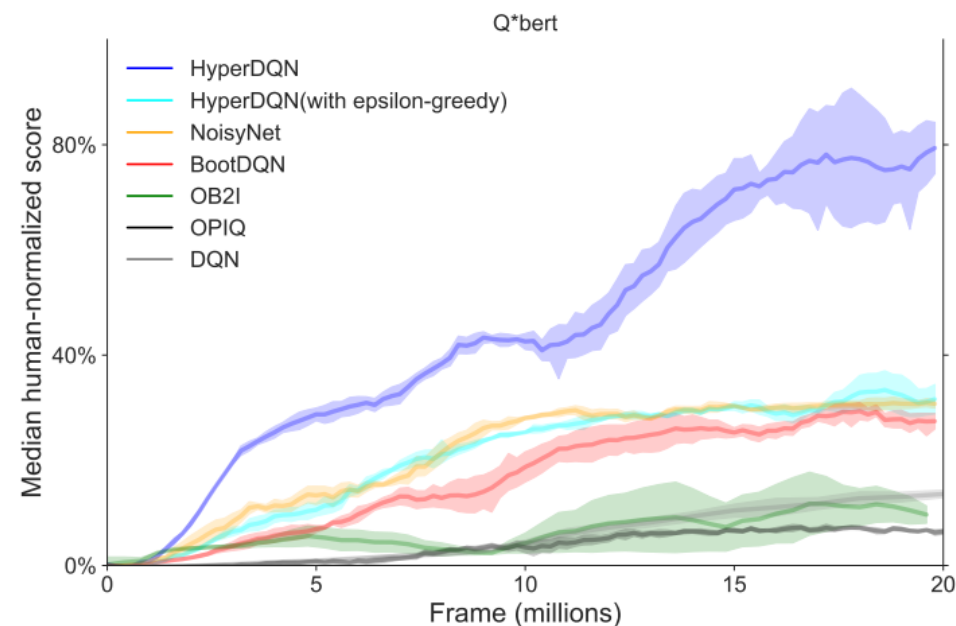
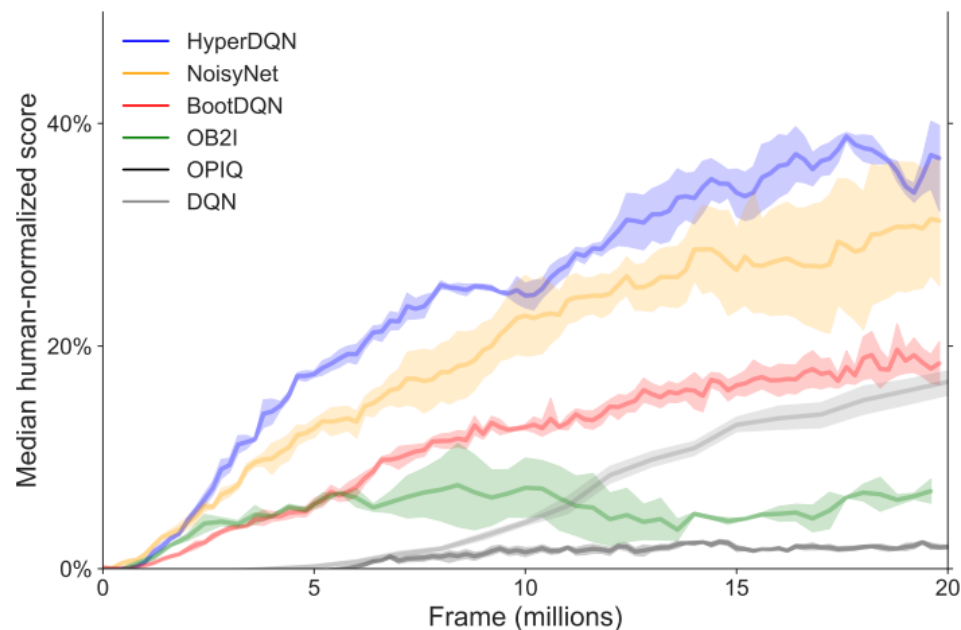


$$\frac{1}{|\tilde{\mathcal{Z}}|} \left(\sum_{z \in \tilde{\mathcal{Z}}} \frac{|\mathcal{D}|}{|\tilde{\mathcal{D}}|} \sum_{(s, a, r, \xi, s') \in \tilde{\mathcal{D}}} \left(Q_{\text{target}}(s', z) + \sigma_\omega z^\top \xi - Q_{\text{prediction}}(s, a, z) \right)^2 + \frac{\sigma_\omega^2}{\sigma_p^2} \|f_\nu(z)\|^2 \right), \quad (\text{A.6})$$

Atari

Table 2: Comparison of algorithms on Atari in terms of the median over 49 games' maximum human-normalized scores. Note that the performance of DQN is based on 200M training frames while other methods are based on 20M training frames.

DQN (200M)	OPIQ	OB2I	BootDQN	NoisyNet	HyperDQN
93%	37%	50%	82%	91%	110%

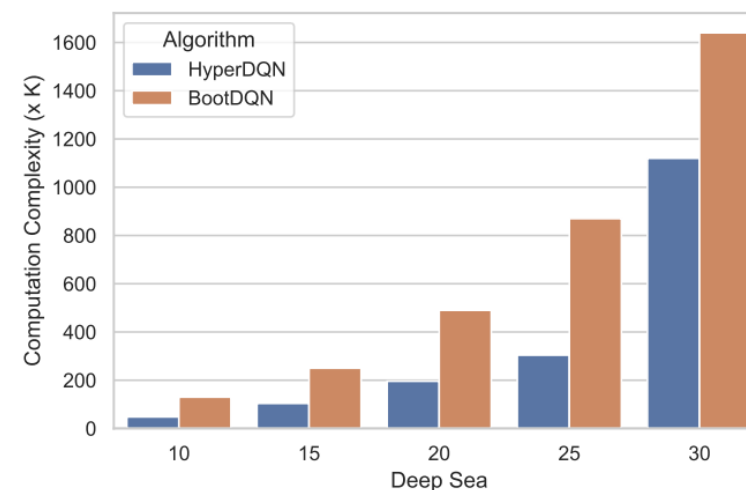
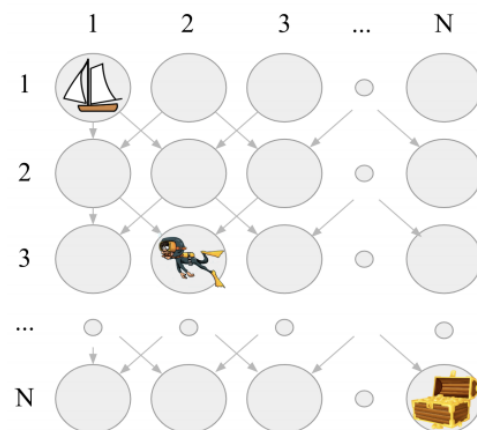


SuperMarioBros

Table 3: Comparison of algorithms on SuperMarioBros in terms of the raw scores by the best policies with 20M training frames.

	DQN	OPIQ	OB2I	BootDQN	NoisyNet	HyperDQN
SuperMarioBros-1-1	1,070	7,650	4,457	7,009	12,439	7,924
SuperMarioBros-1-2	2,883	5,515	4,695	5,665	6,347	8,267
SuperMarioBros-1-3	667	2,053	1,583	1,609	1,587	6,047
SuperMarioBros-2-1	10,800	21,654	14,226	26,415	14,017	23,047
SuperMarioBros-2-2	813	1,630	1,588	1,092	1,808	1,984
SuperMarioBros-2-3	3,373	4,718	4,402	5,108	6,490	5,980
SuperMarioBros-3-1	2,560	3,700	3,251	3,862	11,310	48,385
SuperMarioBros-3-2	11,633	20,872	26,508	20,955	33,489	41,140
SuperMarioBros-3-3	1,007	2,440	3,009	2,650	5,886	5,568

Deep Sea



Thanks