

Reinforcement Learning with Discrete Diffusion Policies

Google Research

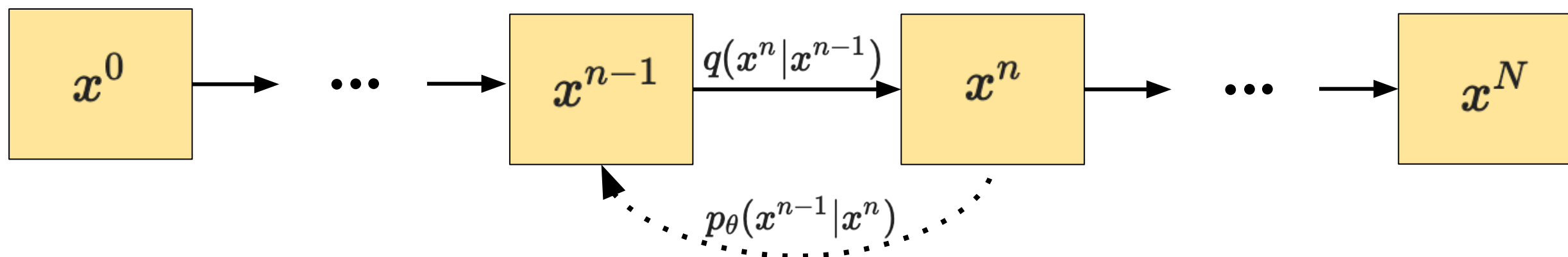
Ofir Nabati

Discrete Diffusion Models

- Clean sample:

$$x^0 = x_{1:K}^0 \in \mathcal{A}^K \quad \mathcal{A} = \{1, 2, 3, \dots, V\}$$

- Forward process: $q(x^n | x^{n-1})$
- (Parameterized) backward process: $p_\theta(x^{n-1} | x^n)$



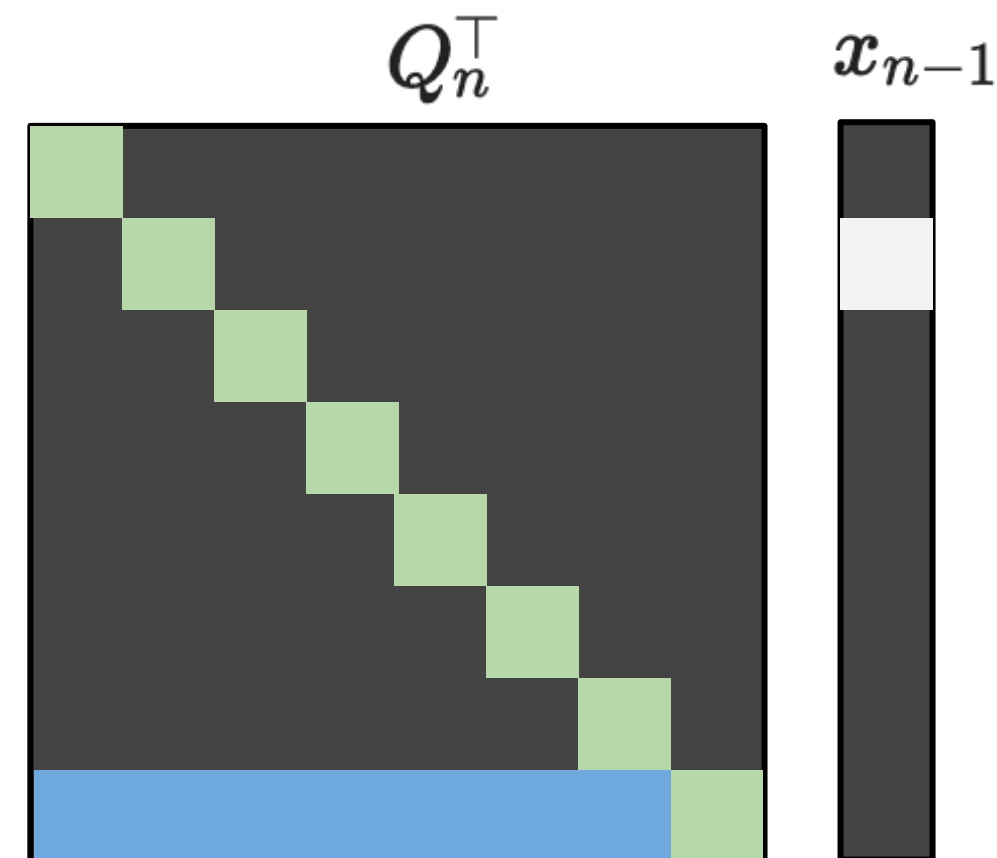
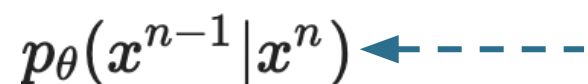
$$q(x^n | x^{n-1}) = \text{Cat}(x^n; Q_n^\top x^{n-1})$$

Q_n is a probability matrix

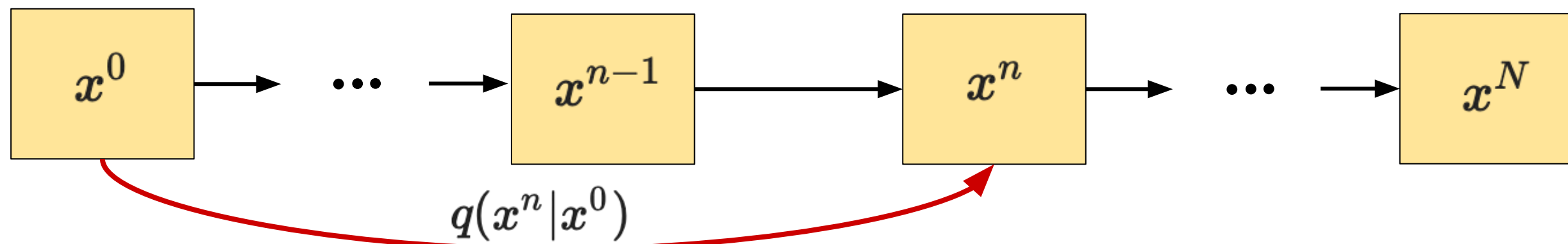
treated as one-hot vector $e_{x^{n-1}}$



$x_n \sim$



Masked Discrete Diffusion Models



The marginal distribution at time n :

$$q(x^n | x^0) = \text{Cat}(x^n; \bar{Q}_n^\top x^0) = \text{Cat}(x^n; \bar{\alpha}_n x^0 + (1 - \bar{\alpha}_n) e_m)$$

$$\bar{Q}_n = \prod_{i=1}^n Q_i = \bar{\alpha}_n I + (1 - \bar{\alpha}_n) \mathbf{1} e_m^\top$$

$$\{\bar{\alpha}_n\}_{n=1}^N \text{ -noise schedule e.g. linear: } \bar{\alpha}_n = 1 - \frac{n}{N}$$

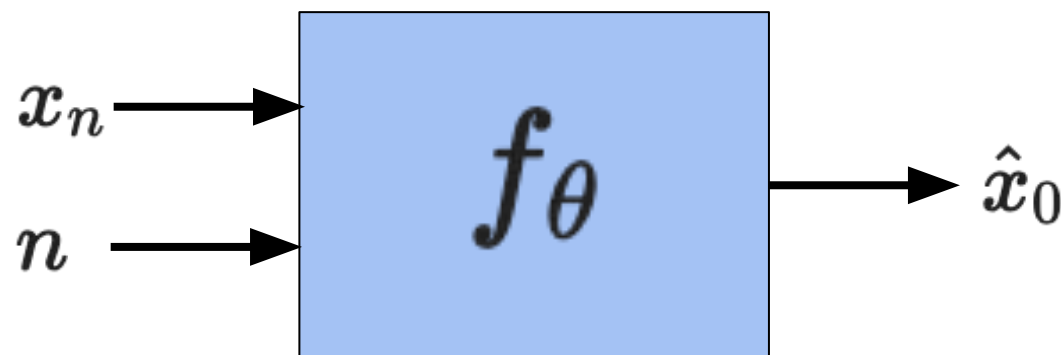
Masked Discrete Diffusion Models

The diffusion objective (ELBO):

$$\log p_{\theta}(x^0) \geq \mathbb{E}_{q(x^1|x^0)} [\log p_{\theta}(x^0|x^1)] - \sum_{n>1} \mathbb{E}_{q(x^n|x^0)} [d_{KL}(\text{posterior } q(x^{n-1}|x^n, x^0) \| p_{\theta}(x^{n-1}|x^n))] - d_{KL}(q(x^N|x^0) \| p(x^N))$$

Usually the backward process is parameterized as the posterior with a **denoiser**

$$p_{\theta}(x^{n-1}|x^n) = q(x^{n-1}|x^n, f_{\theta}(x^n, n))$$



a distribution over values!

Masked Discrete Diffusion Models

For masked discrete diffusion..

$$d_{KL}(q(x^{n-1}|x^n, x^0)||p_\theta(x^{n-1}|x^n)) = -\frac{\bar{\alpha}_{n-1} - \bar{\alpha}_n}{1 - \bar{\alpha}_n} \delta_{x_n, m} \cdot (x^0)^\top \log f_\theta(x^n, n)_{x^0}$$

Estimating the (negative) ELBO

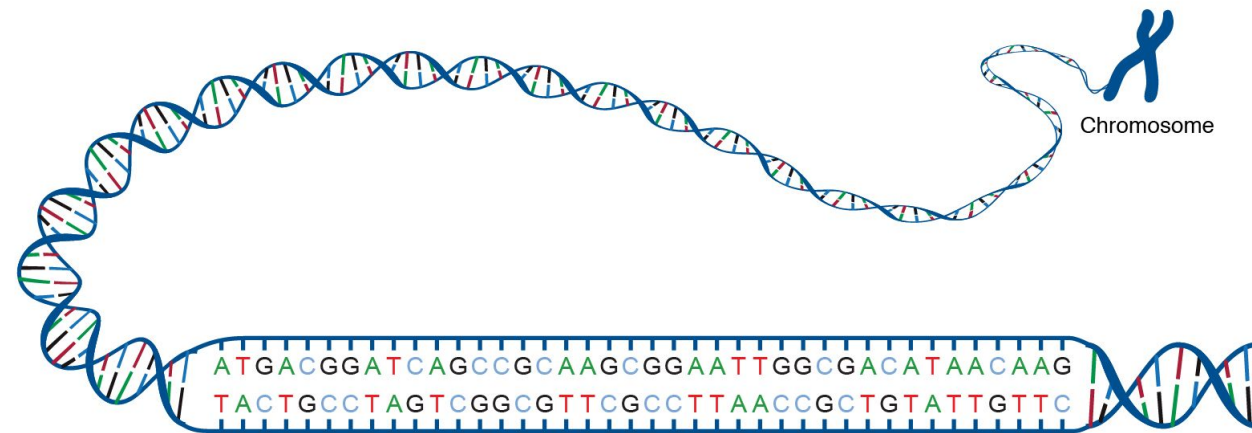
$$\mathcal{L}_{ELBO}(x^0; \theta) \approx \frac{1}{M} \sum_{i=1}^M d_{KL}(q(x^{n_i-1}|x^{n_i}, x^0)||p_\theta(x^{n_i-1}|x^{n_i})) \quad \text{where} \quad n_i \sim \mathcal{U}\{1, 2, \dots, N\}$$

$$x^{n_i} \sim q(x^{n_i}|x^0)$$

Why using discrete diffusion policies?

Combinatorial action spaces:

- Sequence generation
- Action sequences / planning
- Multi agent



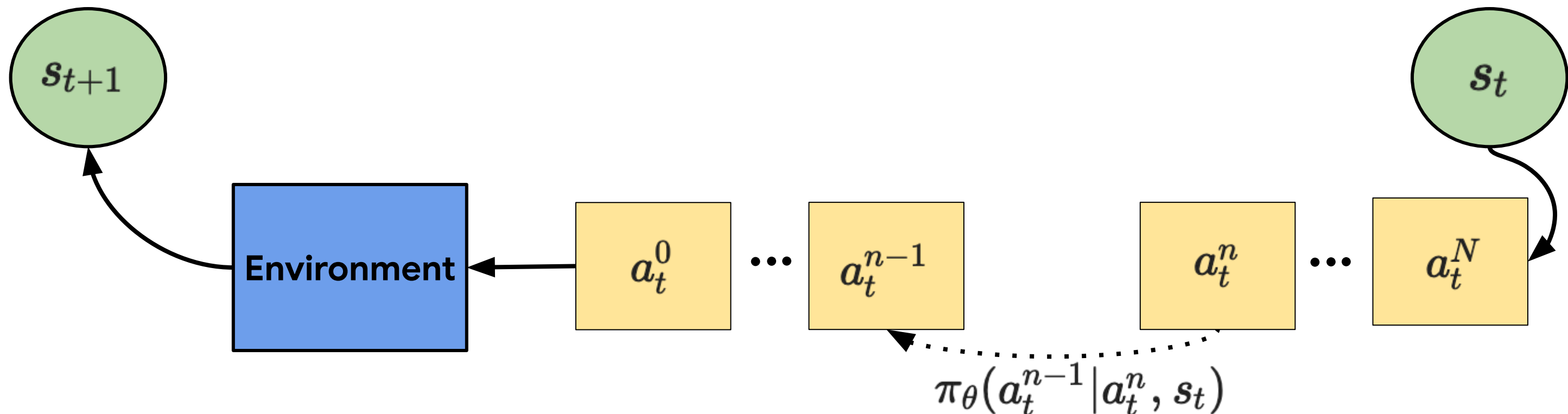
How to model a diffusion policy π_θ ?

Option 1: Augmented MDP

There is only a **backward process**

and it's the **policy** $\pi_\theta(a^{n-1}|a^n, s) = p_\theta(a^{n-1}|a^n, s)$

We create an **augmented MDP** with augmented states: (s_t, a_t^n)

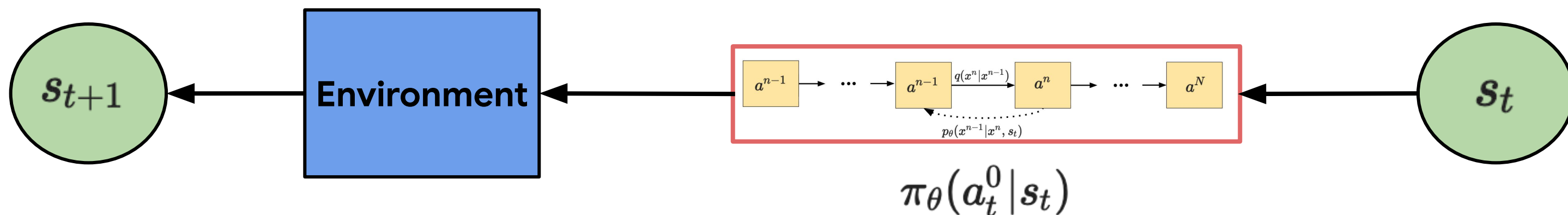


How to model a diffusion policy π_θ ?

Option 2: The diffusion model

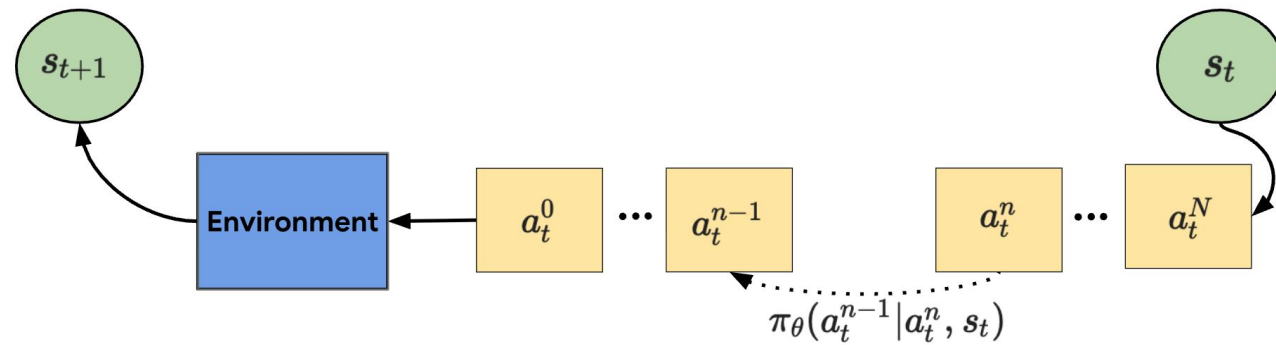
The policy is the **diffusion model**

$$\pi_\theta(a^0 | s) = p_\theta(a^0 | s) = \mathbb{E}_{a^{1:N}} \left[\prod_{n=1}^N p_\theta(a^{n-1} | a^n, s) \right]$$



Which approach is better?

Augmented MDP



Unbiased



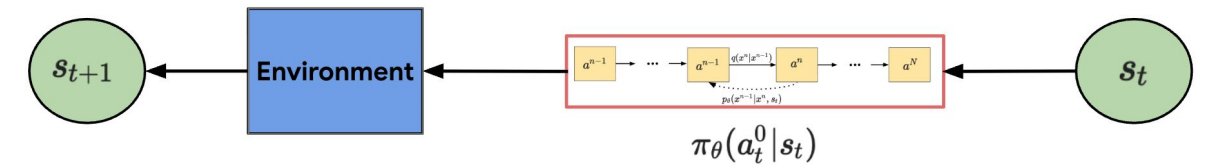
Noisy - high variance



Sample inefficient

e.g: FlowGRPO, DDPO, PPO-variants, DRAKES

Policy as the diffusion model



Biased



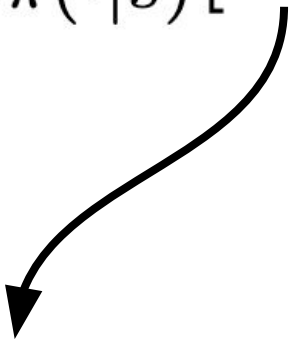
Low variance

- e.g: DiffusionNFT, Self Imitation, BoNC, RWR

Policy Mirror Decent

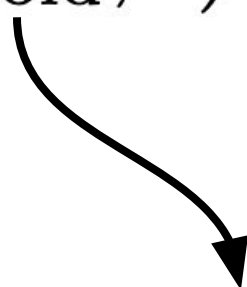
Provably convergent, regularized policy optimization step:

$$\pi(\cdot|s) \in \arg \max_{\pi \in \Pi} \mathbb{E}_{\mathbf{a} \sim \pi(\cdot|s)} [A^{\pi_{\text{old}}}(s, \mathbf{a})] - \lambda d_{KL}(\pi, \pi_{\text{old}}; s) \quad \forall s \in \mathcal{S}$$



A curved arrow points from the $A^{\pi_{\text{old}}}(s, \mathbf{a})$ term in the equation above to the definition below.

$$A^{\pi_{\text{old}}}(s, \mathbf{a}) := q^{\pi_{\text{old}}}(s, \mathbf{a}) - v^{\pi_{\text{old}}}(s)$$



A curved arrow points from the π_{old} term in the equation above to the text below.

previous turn policy

Closed form solution:

$$\pi_{\text{MD}}(\mathbf{a}|s) = \pi_{\text{old}}(\mathbf{a}|s) \exp(A^{\pi_{\text{old}}}(s, \mathbf{a})/\lambda) / Z(s)$$

Policy Mirror Decent

Policy optimization in RL is done in two steps:

Start with initial policy π_0

For $k=1,2,3,\dots$:

- **Policy evaluation:** estimate the Q-function of π_k
- **Policy Improvement:** compute the next-step policy π_{k+1} :

$$\pi_{k+1} \in \arg \min_{\pi_{\theta} \in \Pi} \mathbb{E}_{s \sim \mathcal{D}} [d(\pi_{\theta}, \pi_{\text{MD}}^k; s)]$$

$$\pi_{\text{MD}}(\mathbf{a}|s) = \pi_{\text{old}}(\mathbf{a}|s) \exp(A^{\pi_{\text{old}}}(s, \mathbf{a})/\lambda) / Z(s)$$

Forward KL+ Diffusion Policy (Self Imitation)

The choice of the distance determine the objective:

Forward KL: $d_{KL}(\pi_{MD}^k || \pi_{\theta}; s)$

+ policy as the diffusion model: $-\log \pi_{\theta}(a^0 | s) \leq \mathcal{L}_{ELBO}(a^0, s; \theta)$

$$\mathcal{L}_{FKL}(\theta) = \mathbb{E}_{s \sim \mathcal{D}, \hat{\mathcal{A}}_s \sim \pi_k} \left[\sum_{\mathbf{a}^0 \in \hat{\mathcal{A}}_s} (\text{softmax}_{\mathbf{a} \in \hat{\mathcal{A}}_s} (A^{\pi_k}(s, \mathbf{a}^0) / \lambda)) \cdot \mathcal{L}_{ELBO}(\mathbf{a}^0, s; \theta) \right]$$

sampled action set
Softmax weights
temperature

Reverse KL+ Diffusion Policy

Reverse KL: $d_{KL}(\pi_\theta || \pi_{MD}^k; s)$

+ policy as the backward process (augmented MDP):

$$-\log \pi_\theta(a^{n-1} | a^n, s) = -\log p_\theta(a^{n-1} | a^n, s)$$

$$\mathcal{L}_{\text{RKL}}(\theta) = \mathbb{E}_{s \sim \mathcal{D}, \mathbf{a} \sim \pi_k} [-\eta(s, \mathbf{a}; \theta) A^{\pi_k}(s, \mathbf{a}) + \lambda d_{KL}(\pi_\theta, \pi_k; s)]$$

Importance sampling: $\eta((s, \mathbf{a}^n), \mathbf{a}^{n-1}; \theta) = \frac{p_\theta(\mathbf{a}^{n-1} | \mathbf{a}^n, s)}{p_k(\mathbf{a}^{n-1} | \mathbf{a}^n, s)}$

Experiments

Goals

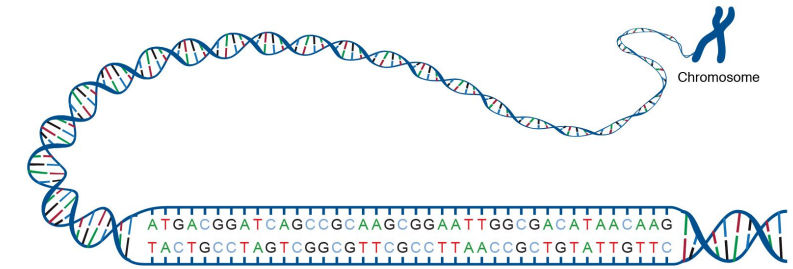
- Better performance (high return)
- Better efficiency compared to AR:

Number of diffusion steps (N) < sequence length (T)

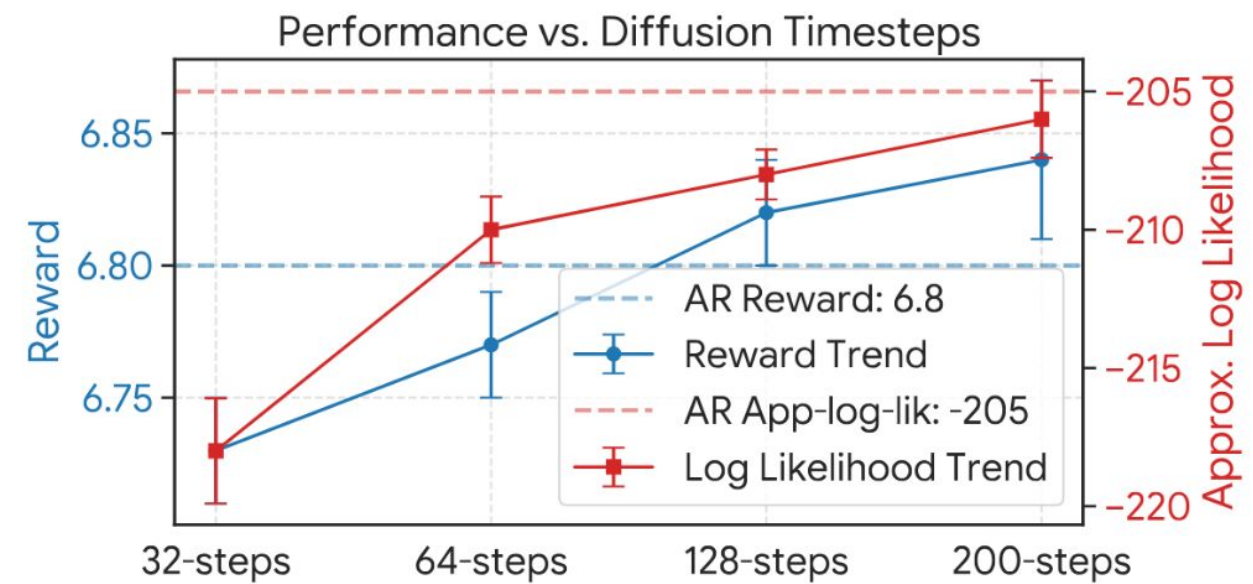
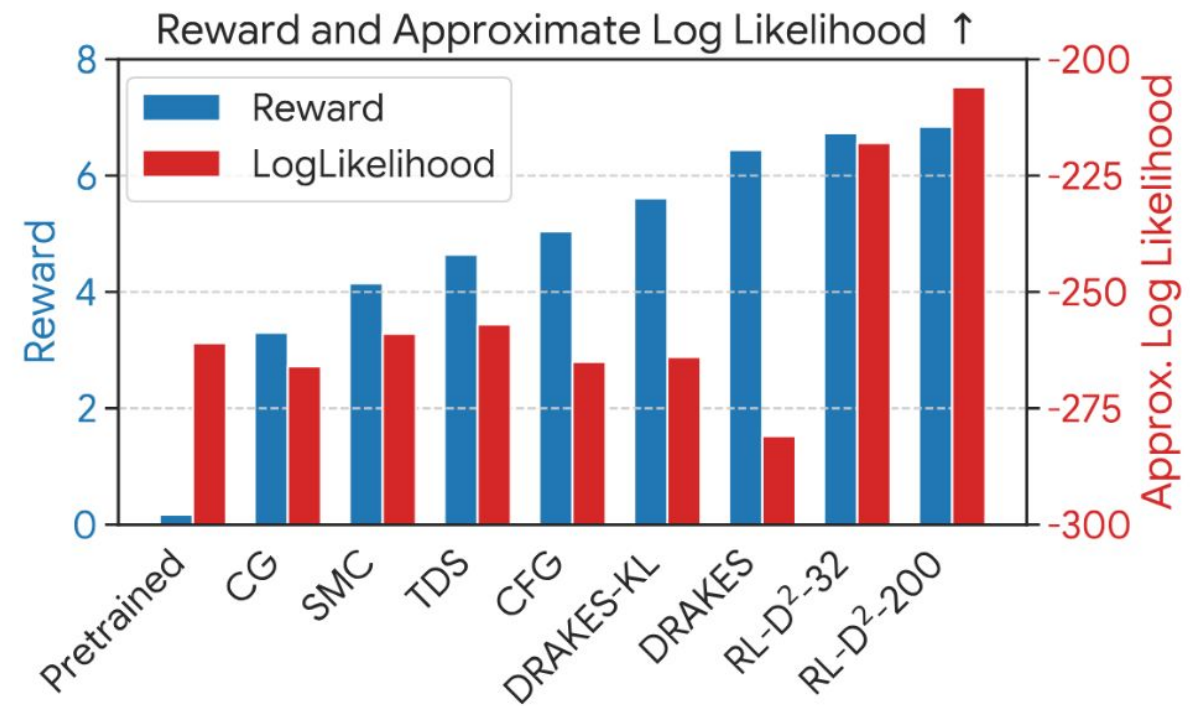
Experiments

DNA sequence generation (single step)

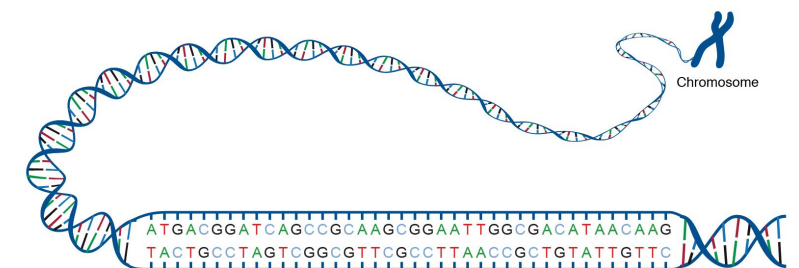
- RL fine tuning w.r.t a pretrained model
- Reward measure **Predicted Activity** in liver cancer cells
- Sequence length: 200
- Action space size: 4^{200}
- Objective: Self Imitation (FKL)



Experiments

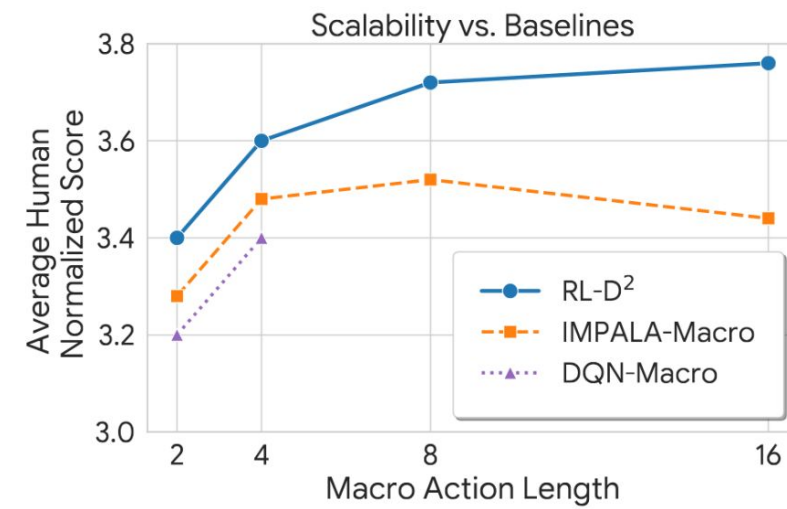
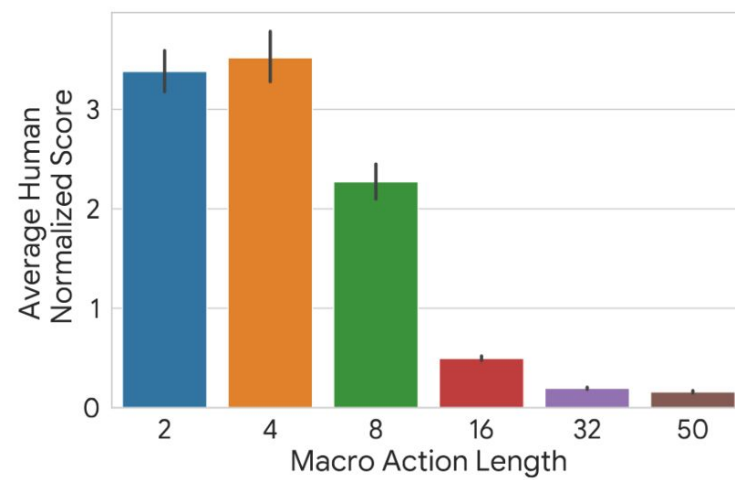
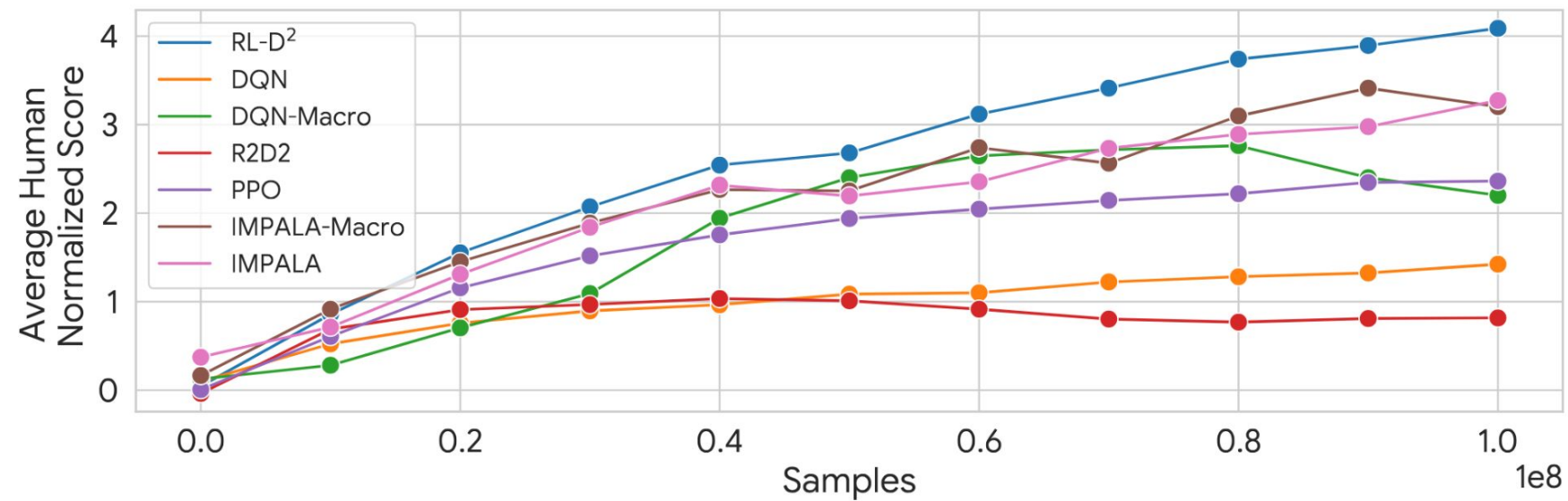


- Outperform AR with half the number of steps
- SOTA compared to DRAKES



Experiments

Action sequence generation (Atari)



Experiments

Google Research Football

- From scratch
- Different scenarios: 11vs11, 5vs5, corner
- Sequence length: players number
- Action space size: $19^{\#players}$
- Objective: RKL (PPO) and FKL
- Reward: +1 for goal, -1 for the opponent goal
- Stochastic environment



Experiments

	PASS & SHOOT	RUN PASS & SHOT	3 VS 1	CT-EASY	CT-HARD	CORNER	5 VS 5	11 VS 11	small budget 11 VS 11 SM
AUTOREGRESSIVE	78.2 $\pm$ 40.1	70.2 $\pm$ 35.2	75.8 $\pm$ 35.1	69.3 $\pm$ 30.6	33.6 $\pm$ 40.1	0.0 $\pm$ 0.0	71.1 $\pm$ 30.5	34.4 $\pm$ 14.4	0.0 $\pm$ 0.0
RL-D ² RKL (OURS)	100 $\pm$ 0.0	98.5 $\pm$ 1.6	98.6 $\pm$ 1.0	99.1 $\pm$ 0.6	97.7 $\pm$ 2.2	96.3 $\pm$ 3.4	99.6 $\pm$ 0.4	99.8 $\pm$ 0.2	0.0 $\pm$ 0.0
RL-D ² FKL (OURS)	99.0 $\pm$ 1.0	98.6 $\pm$ 0.5	97.3 $\pm$ 2.0	98.0 $\pm$ 1.1	97.1 $\pm$ 2.9	93.2 $\pm$ 2.9	98.4 $\pm$ 1.0	97.4 $\pm$ 0.9	67.2 $\pm$ 8.1
MAT	97.9 $\pm$ 2.1	98.3 $\pm$ 1.2	92.9 $\pm$ 1.1	87.9 $\pm$ 2.1	88.2 $\pm$ 3.8	95.3 $\pm$ 2.5	93.7 $\pm$ 0.9	92.0 $\pm$ 2.4	9.0 $\pm$ 3.7
MAPPO	99.5 $\pm$ 0.2	73.2 $\pm$ 3.6	93.2 $\pm$ 1.5	70.1 $\pm$ 3.8	63 $\pm$ 2.1	53.1 $\pm$ 5.3	95.4 $\pm$ 1.6	52.6 $\pm$ 3.1	5.0 $\pm$ 0.2

- Practically solve Gfootball ~100% success for all scenarios
- Outperform SOTA AR method (MAT)
- FKL much more efficient but asymptotically underperform
- Needs a lot of samples..