

S²AC: ENERGY-BASED REINFORCEMENT LEARNING WITH STEIN SOFT ACTOR CRITIC

Safa Messaoud^{1†}, Billel Mokeddem^{1*}, Zhenghai Xue^{2*}, Linsey Pang³, Bo An^{4,2},

Haipeng Chen^{5†}, Sanjay Chawla^{1†}

¹Qatar Computing Research Institute, Hamad Bin Khalifa University, ²School of Computer Science and Engineering, Nanyang Technological University, ³SalesForce, ⁴Skywork AI, ⁵Data Science, William & Mary
{smessaoud, bmokeddem, schawla}@hbku.edu.qa, zhenghai001@e.ntu.edu.sg
panglinsey@gmail.com, boan@ntu.edu.sg, hchen23@wm.edu

* Equal contribution † Corresponding authors

ABSTRACT

Learning expressive stochastic policies instead of deterministic ones has been proposed to achieve better stability, sample complexity, and robustness. Notably, in Maximum Entropy Reinforcement Learning (MaxEnt RL), the policy is modeled as an expressive Energy-Based Model (EBM) over the Q-values. However, this formulation requires the estimation of the entropy of such EBMs, which is an open problem. To address this, previous MaxEnt RL methods either implicitly estimate the entropy, resulting in high computational complexity and variance (SQL), or follow a variational inference procedure that fits simplified actor distributions (*e.g.*, Gaussian) for tractability (SAC). We propose Stein Soft Actor-Critic (S²AC), a MaxEnt RL algorithm that learns expressive policies without compromising efficiency. Specifically, S²AC uses parameterized Stein Variational Gradient Descent (SVGD) as the underlying policy. We derive a closed-form expression of the entropy of such policies. Our formula is computationally efficient and only depends on first-order derivatives and vector products. Empirical results show that S²AC yields more optimal solutions to the MaxEnt objective than SQL and SAC in the multi-goal environment, and outperforms SAC and SQL on the MuJoCo benchmark. Our code is available at: <https://github.com/SafaMessaoud/S2AC-Energy-Based-RL-with-Stein-Soft-Actor-Critic>

1 INTRODUCTION

MaxEnt RL (Todorov, 2006; Ziebart, 2010; Haarnoja et al., 2017; Kappen, 2005; Toussaint, 2009; Theodorou et al., 2010; Abdolmaleki et al., 2018; Haarnoja et al., 2018a; Vieillard et al., 2020) has been proposed to address challenges hampering the deployment of RL to real-world applications, including stability, sample efficiency (Gu et al., 2017), and robustness (Eysenbach & Levine, 2022). Instead of learning a deterministic policy, as in classical RL (Sutton et al., 1999; Schulman et al., 2017; Silver et al., 2014; Lillicrap et al., 2015), MaxEnt RL learns a stochastic policy that captures the intricacies of the action space. This enables better exploration during training and eventually better robustness to environmental perturbations at test time, *i.e.*, the agent learns multimodal action space distributions which enables picking the next best action in case a perturbation prevents the execution of the optimal one. To achieve this, MaxEnt RL models the policy using the expressive family of EBMs (LeCun et al., 2006). This translates into learning policies that maximize the sum of expected future reward and expected future entropy. However, estimating the entropy of such complex distributions remains an open problem.

To address this, existing approaches either use tricks to go around the entropy computation or make limiting assumptions on the policy. This results in either poor scalability or convergence to suboptimal solutions. For example, SQL (Haarnoja et al., 2017) implicitly incorporates entropy in the Q-function computation. This requires using importance sampling, which results in high variability and hence poor training stability and limited scalability to high dimensional action spaces. SAC (Haarnoja

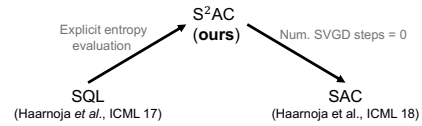


Figure 1: Comparing S²AC to SQL and SAC. S²AC with a parameterized policy is reduced to SAC if the number of SVGD steps is 0. SQL becomes equivalent to S²AC if the entropy is evaluated explicitly with our derived formula.

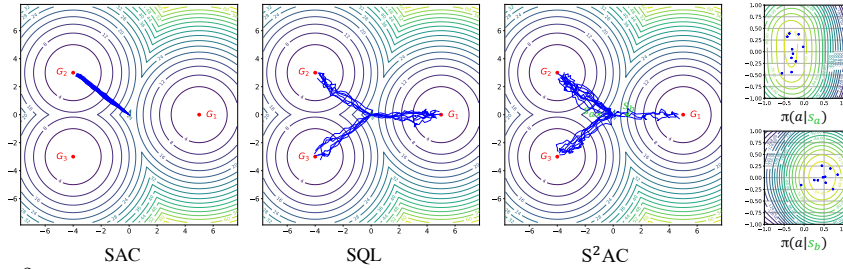


Figure 2: S^2AC learns a more optimal solution to the MaxEnt RL objective than SAC and SQL. We design a multi-goal environment where an agent starts from the center of the 2-d map and tries to reach one of the three goals (G_1 , G_2 , and G_3). The maximum expected future reward (level curves) is the same for all the goals but the expected future entropy is different (higher on the path to G_2/G_3): the action distribution $\pi(a|s)$ is bi-modal on the path to the left (G_2 and G_3) and unimodal to the right (G_1). Hence, we expect the optimal policy for the MaxEnt RL objective to assign more weights to G_2 and G_3 . We visualize trajectories (in blue) sampled from the policies learned using SAC, SQL, and S^2AC . SAC quickly commits to a single mode due to its actor being tied to a Gaussian policy. Though SQL also recovers the three modes, the trajectories are evenly distributed. S^2AC recovers all the modes and approaches the left two goals more frequently. This indicates that it successfully maximizes not only the expected future reward but also the expected future entropy.

et al., 2018a), on the other hand, follows a variational inference procedure by fitting a Gaussian distribution to the EBM policy. This enables a closed-form evaluation of the entropy but results in a suboptimal solution. For instance, SAC fails in environments characterized by multimodal action distributions. Similar to SAC, IAPO (Marino et al., 2021) models the policy as a uni-modal Gaussian. Instead of optimizing a MaxEnt objective, it achieves multimodal policies by learning a collection of parameter estimates (mean, variance) through different initializations for different policies. To improve the expressiveness of SAC, SSPG (Cetin & Celiktutan, 2022) and SAC-NF (Mazouze et al., 2020) model the policy as a Markov chain with Gaussian transition probabilities and as a normalizing flow (Rezende & Mohamed, 2015), respectively. However, due to training stability issues, the reported results in Cetin & Celiktutan (2022) show that though both models learn multi-modal policies, they fail to maximize the expected future entropy in positive rewards setups.

We propose a new algorithm, S^2AC , that yields a more optimal solution to the MaxEnt RL objective. To achieve *expressivity*, S^2AC models the policy as a Stein Variational Gradient Descent (SVGD) (Liu, 2017) sampler from an EBM over Q-values (target distribution). SVGD proceeds by first sampling a set of particles from an initial distribution, and then iteratively transforming these particles via a sequence of updates to fit the target distribution. To compute a *closed-form estimate of the entropy* of such policies, we use the change-of-variable formula for pdfs (Devore et al., 2012). We prove that this is only possible due to the invertibility of the SVGD update rule, which does not necessarily hold for other popular samplers (e.g., Langevin Dynamics (Welling & Teh, 2011)). While normalizing flow models (Rezende & Mohamed, 2015) are also invertible, SVGD-based policy is more expressive as it encodes the inductive bias about the unnormalized density and incorporates a dispersion term to encourage multi-modality, whereas normalizing flows encode a restrictive class of invertible transformations (with easy-to-estimate Jacobian determinants). Moreover, our formula is computationally efficient and only requires evaluating first-order derivatives and vector products. To improve *scalability*, we model the initial distribution of the SVGD sampler as an isotropic Gaussian and learn its parameters, *i.e.*, mean and standard deviation, end-to-end. We show that this results in faster convergence to the target distribution, *i.e.*, fewer SVGD steps. Intuitively, the initial distribution learns to contour the high-density region of the target distribution while the SVGD updates result in better and faster convergence to the modes within that region. Hence, our approach is as parameter efficient as SAC, since the SVGD updates do not introduce additional trainable parameters.

Note that S^2AC can be reduced to SAC when the number of SVGD steps is zero. Also, SQL becomes equivalent to S^2AC if the entropy is computed explicitly using our formula (the policy in SQL is an amortized SVGD sampler). Beyond RL, the backbone of S^2AC is a new variational inference algorithm with a more expressive and scalable distribution characterized by a closed-form entropy estimate. We believe that this variational distribution can have a wider range of exciting applications.

We conduct extensive empirical evaluations of S^2AC from three aspects. We start with a sanity check on the merit of our derived SVGD-based entropy estimate on target distributions with known entropy values (e.g., Gaussian) or log-likelihoods (e.g., Gaussian Mixture Models) and assess its

sensitivity to different SVGD parameters (kernel, initial distribution, number of steps and number of particles). We observe that its performance depends on the choice of the kernel and is robust to variations of the remaining parameters. In particular, we find out that the kernel should be chosen to guarantee inter-dependencies between the particles, which turns out to be essential for invertibility. Next, we assess the performance of S^2AC on a multi-goal environment (Haarnoja et al., 2017) where different goals are associated with the same positive (maximum) expected future reward but different (maximum) expected future entropy. We show that S^2AC learns multimodal policies and effectively maximizes the entropy, leading to better robustness to obstacles placed at test time. Finally, we test S^2AC on the MuJoCo benchmark (Duan et al., 2016). S^2AC yields better performances than the baselines on four out of the five environments. Moreover, S^2AC shows higher sample efficiency as it tends to converge with fewer training steps. These results were obtained from running SVGD for only three steps, which results in a small overhead compared to SAC during training. Furthermore, to maximize the run-time efficiency during testing, we train an amortized SVGD version of the policy to mimic the SVGD-based policy. Hence, this reduces inference to a forward pass through the policy network without compromising the performance.

2 PRELIMINARIES

2.1 SAMPLERS FOR ENERGY-BASED MODELS

In this work, we study three representative methods for sampling from EBMs: (1) Stochastic Gradient Langevin Dynamics (SGLD) & Deterministic Langevin Dynamics (DLD) (Welling & Teh, 2011), (2) Hamiltonian Monte Carlo (HMC) (Neal et al., 2011), and (3) Stein Variational Gradient Descent (SVGD) (Liu & Wang, 2016). We review SVGD here since it is the sampler we eventually use in S^2AC , and leave the rest to Appendix C.1. SVGD is a particle-based Bayesian inference algorithm. Compared to SGLD and HMC which have a single particle in their dynamics, SVGD operates on a set of particles. Specifically, SVGD samples a set of m particles $\{a_j\}_{j=1}^m$ from an initial distribution q^0 which it then transforms through a sequence of updates to fit the target distribution. Formally, at every iteration l , SVGD applies a form of functional gradient descent Δf that minimizes the KL-divergence between the target distribution p and the proposal distribution q^l induced by the particles, *i.e.*, the update rule for the i^{th} particles is: $a_i^{l+1} = a_i^l + \epsilon \Delta f(a_i^l)$ with

$$\Delta f(a_i^l) = \mathbb{E}_{a_j \sim q^l} [k(a_i^l, a_j^l) \nabla_{a_j^l} \log p(a_j^l) + \nabla_{a_j^l} k(a_i^l, a_j^l)]. \quad (1)$$

Here, ϵ is the step size and $k(\cdot, \cdot)$ is the kernel function, *e.g.*, the RBF kernel: $k(a_i, a_j) = \exp(-\|a_i - a_j\|^2 / 2\sigma^2)$. The first term within the gradient drives the particles toward the high probability regions of p , while the second term serves as a repulsive force to encourage dispersion.

2.2 MAXIMUM-ENTROPY RL

We consider an infinite horizon Markov Decision Process (MDP) defined by a tuple $(\mathcal{S}, \mathcal{A}, p, r)$, where $\mathcal{S}$ is the state space, $\mathcal{A}$ is the action space and $p : \mathcal{S} \times \mathcal{A} \times \mathcal{S} \rightarrow [0, \infty]$ is the state transition probability modeling the density of the next state $s_{t+1} \in \mathcal{S}$ given the current state $s_t \in \mathcal{S}$ and action $a_t \in \mathcal{A}$. Additionally, we assume that the environment emits a bounded reward function $r \in [r_{\min}, r_{\max}]$ at every iteration. We use $\rho_\pi(s_t)$ and $\rho_\pi(s_t, a_t)$ to denote the state and state-action marginals of the trajectory distribution induced by a policy $\pi(a_t|s_t)$. We consider the setup of continuous action spaces (Lazaric et al., 2007; Lee et al., 2018; Zhou & Lu, 2023). MaxEnt RL (Todorov, 2006; Ziebart, 2010; Rawlik et al., 2012) learns a policy $\pi^*(a_t|s_t)$, that instead of maximizing the expected future reward, maximizes the sum of the expected future reward and entropy:

$$\pi^* = \arg \max_{\pi} \sum_t \gamma^t \mathbb{E}_{(s_t, a_t) \sim \rho_\pi} [r(s_t, a_t) + \alpha \mathcal{H}(\pi(\cdot|s_t))], \quad (2)$$

where α is a temperature parameter controlling the stochasticity of the policy and $\mathcal{H}(\pi(\cdot|s_t))$ is the entropy of the policy at state s_t . The conventional RL objective can be recovered for $\alpha = 0$. Note that the MaxEnt RL objective above is equivalent to approximating the policy, modeled as an EBM over Q-values, by a variational distribution $\pi(a_t|s_t)$ (see proof of equivalence in Appendix D), *i.e.*,

$$\pi^* = \arg \min_{\pi} \sum_t \mathbb{E}_{s_t \sim \rho_\pi} [D_{KL}(\pi(\cdot|s_t) \| \exp(Q(s_t, \cdot)/\alpha)/Z)], \quad (3)$$

where D_{KL} is the KL-divergence and Z is the normalizing constant. We now review two landmark MaxEnt RL algorithms: SAC (Haarnoja et al., 2018a) and SQL (Haarnoja et al., 2017).

SAC is an actor-critic algorithm that alternates between policy evaluation, *i.e.*, evaluating the Q-values for a policy $\pi_\theta(a_t|s_t)$:

$$Q_\phi(s_t, a_t) \leftarrow r(s_t, a_t) + \gamma \mathbb{E}_{s_{t+1}, a_{t+1} \sim \rho_{\pi_\theta}} [Q_\phi(s_{t+1}, a_{t+1}) + \alpha \mathcal{H}(\pi_\theta(\cdot|s_{t+1}))] \quad (4)$$

and policy improvement, *i.e.*, using the updated Q-values to compute a better policy:

$$\pi_\theta = \arg \max_\theta \sum_t \mathbb{E}_{s_t, a_t \sim \rho_{\pi_\theta}} [Q_\phi(a_t, s_t) + \alpha \mathcal{H}(\pi_\theta(\cdot|s_t))]. \quad (5)$$

SAC models π_θ as an isotropic Gaussian, *i.e.*, $\pi_\theta(\cdot|s) = \mathcal{N}(\mu_\theta, \sigma_\theta I)$. While this enables computing a closed-form expression of the entropy, it incurs an over-simplification of the true action distribution, and thus cannot represent complex distributions, *e.g.*, multimodal distributions.

SQL goes around the entropy computation, by defining a soft version of the value function $V_\phi = \alpha \log(\int_{\mathcal{A}} \exp(\frac{1}{\alpha} Q_\phi(s_t, a')) da')$. This enables expressing the Q-value (Eq (4)) independently from the entropy, *i.e.*, $Q_\phi(s_t, a_t) = r(s_t, a_t) + \gamma \mathbb{E}_{s_{t+1} \sim p}[V_\phi(s_{t+1})]$. Hence, SQL follows a soft value iteration which alternates between the updates of the “soft” versions of Q and value functions:

$$Q_\phi(s_t, a_t) \leftarrow r(s_t, a_t) + \gamma \mathbb{E}_{s_{t+1} \sim p}[V_\phi(s_{t+1})], \quad \forall (s_t, a_t) \quad (6)$$

$$V_\phi(s_t) \leftarrow \alpha \log(\int_{\mathcal{A}} \exp(\frac{1}{\alpha} Q_\phi(s_t, a')) da'), \quad \forall s_t. \quad (7)$$

Once the Q_ϕ and V_ϕ functions converge, SQL uses amortized SVGD [Wang & Liu \(2016\)](#) to learn a stochastic sampling network $f_\theta(\xi, s_t)$ that maps noise samples ξ into the action samples from the EBM policy distribution $\pi^*(a_t|s_t) = \exp(\frac{1}{\alpha}(Q^*(s_t, a_t) - V^*(s_t)))$. The parameters θ are obtained by minimizing the loss $J_\theta(s_t) = D_{KL}(\pi_\theta(\cdot|s_t) || \exp(\frac{1}{\alpha}(Q^*(s_t, \cdot) - V^*(s_t))))$ with respect to θ . Here, π_θ denotes the policy induced by f_θ . SVGD is designed to minimize such KL-divergence without explicitly computing π_θ . In particular, SVGD provides the most greedy direction as a functional $\Delta f_\theta(\cdot, s_t)$ (Eq (1)) which can be used to approximate the gradient $\partial J_\theta / \partial a_t$. Hence, the gradient of the loss J_θ with respect to θ is: $\partial J_\theta(s_t) / \partial \theta \propto \mathbb{E}_\xi [\Delta f_\theta(\xi, s_t) \partial f_\theta(\xi, s_t) / \partial \theta]$. Note that the integral in Eq (7) is approximated via importance sampling, which is known to result in high variance estimates and hence poor scalability to high dimensional action spaces. Moreover, amortized generation is usually unstable and prone to mode collapse, an issue similar to GANs. Therefore, SQL is outperformed by SAC [Haarnoja et al. \(2018a\)](#) on benchmark tasks like MuJoCo.

3 APPROACH

We introduce S²AC, a new actor-critic MaxEnt RL algorithm that uses SVGD as the underlying actor to generate action samples from policies represented using EBMs. This choice is motivated by the expressivity of distributions that can be fitted via SVGD. Additionally, we show that we can derive a closed-form entropy estimate of the SVGD-induced distribution, thanks to the invertibility of the update rule, which does not necessarily hold for other EBM samplers. Besides, we propose a parameterized version of SVGD to enable scalability to high-dimensional action spaces and non-smooth Q-function landscapes. S²AC is hence capable of learning a more optimal solution to the MaxEnt RL objective (Eq (2)) as illustrated in Figure 2.

3.1 STEIN SOFT ACTOR CRITIC

Like SAC, S²AC performs soft policy iteration which alternates between policy evaluation and policy improvement. The difference is that we model the actor as a *parameterized sampler from an EBM*. Hence, the policy distribution corresponds to an expressive EBM as opposed to a Gaussian.

Critic. The critic’s parameters ϕ are obtained by minimizing the Bellman loss as traditionally:

$$\phi^* = \arg \min_\phi \mathbb{E}_{(s_t, a_t) \sim \rho_{\pi_\theta}} [(Q_\phi(s_t, a_t) - \hat{y})^2], \quad (8)$$

with the target $\hat{y} = r_t(s_t, a_t) + \gamma \mathbb{E}_{(s_{t+1}, a_{t+1}) \sim \rho_\pi} [Q_\phi(s_{t+1}, a_{t+1}) + \alpha \mathcal{H}(\pi(\cdot|s_{t+1}))]$. Here $\bar{\phi}$ is an exponentially moving average of the value network weights [\(Mnih et al. 2015\)](#).

Actor as an EBM sampler. The actor is modeled as a sampler from an EBM over the Q-values. To generate a set of valid actions, the actor first samples a set of particles $\{a^0\}$ from an initial distribution q^0 (*e.g.*, Gaussian). These particles are then updated over several iterations $l \in [1, L]$, *i.e.*, $\{a^{l+1}\} \leftarrow \{a^l\} + \epsilon h(\{a^l\}, s)$ following the sampler dynamics characterized by a transformation h (*e.g.*, for SVGD, $h = \Delta f$ in Eq (1)). If q^0 is tractable and h is invertible, it’s possible to compute a closed-form expression of the distribution of the particles at the l^{th} iteration via the change of variable formula [Devore et al. \(2012\)](#): $q^l(a^l|s) = q^{l-1}(a^{l-1}|s) |\det(I + \epsilon \nabla_{a^l} h(a^l, s))|^{-1}$, $\forall l \in [1, L]$. In this case, the policy is represented using the particle distribution at the final step L of the sampler dynamics, *i.e.*, $\pi(a|s) = q^L(a^L|s)$ and the entropy can be estimated by averaging $\log q^L(a^L|s)$ over a set of particles (Section 3.2). We study the invertibility of popular EBM samplers in Section 3.3.

Parameterized initialization. To reduce the number of steps required to converge to the target distribution (hence reducing computation cost), we further propose modeling the initial distribution as a parameterized isotropic Gaussian, *i.e.*, $a^0 \sim \mathcal{N}(\mu_\theta(s), \sigma_\theta(s))$. The parameterization trick is then used to express a^0 as a function of θ . Intuitively, the actor would learn θ such that the initial distribution is close to the target distribution. Hence, fewer steps are required to converge, as illustrated in Figure 3. Note that if the number of steps $L = 0$, S^2AC is reduced to SAC. Besides, to deal with the non-smooth nature of deep Q-function landscapes which might lead to particle divergence in the sampling process, we bound the particle updates to be within a few standard deviations (t) from the mean of the learned initial distribution, *i.e.*, $-\tau\sigma_\theta \leq a_\theta^l \leq \tau\sigma_\theta$, $\forall l \in [1, L]$. Eventually, the initial distribution q_θ^0 learns to contour the high-density region of the target distribution and the following updates refine it by converging to the spanned modes. Formally, the parameters θ are computed by minimizing the expected KL-divergence between the policy q_θ^L induced by the particles from the sampler and the EBM of the Q-values:

$$\begin{aligned} \theta^* = \arg \max_{\theta} \mathbb{E}_{s_t \sim \mathcal{D}, a_\theta^L \sim \pi_\theta} [Q_\phi(s_t, a_\theta^L)] + \alpha \mathbb{E}_{s_t \sim \mathcal{D}} [\mathcal{H}(\pi_\theta(\cdot|s_t))] \\ \text{s.t. } -t\sigma_\theta \leq a_\theta^l \leq t\sigma_\theta, \quad \forall l \in [1, L]. \end{aligned} \quad (9)$$

Here, $\mathcal{D}$ is the replay buffer. The derivation is in Appendix E. Note that the constraint does not truncate the particles as it is not an invertible transformation which then violates the assumptions of the change of variable formula. Instead, we sample more particles than we need and select the ones that stay within the range. We call $S^2AC(\phi, \theta)$ and $S^2AC(\phi)$ as two versions of S^2AC with/without the parameterized initial distribution. The complete S^2AC algorithm is in Algorithm 1 of Appendix A.

3.2 A CLOSED-FORM EXPRESSION OF THE POLICY’S ENTROPY

A critical challenge in MaxEnt RL is how to efficiently compute the entropy term $\mathcal{H}(\pi(\cdot|s_{t+1}))$ in Eq (2). We show that, if we model the policy as an iterative sampler from the EBM, under certain conditions, we can derive a closed-form estimate of the entropy at convergence.

Theorem 3.1. *Let $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be an invertible transformation of the form $F(a) = a + \epsilon h(a)$. We denote by $q^L(a^L)$ the distribution obtained from repeatedly applying F to a set of samples $\{a^0\}$ from an initial distribution $q^0(a^0)$ over L steps, *i.e.*, $a^L = F \circ F \circ \dots \circ F(a^0)$. Under the condition $\epsilon \|\nabla_{a_i} h(a_i)\|_\infty \ll 1, \forall i \in [1, L]$, the distribution of the particles at the L^{th} step is:*

$$\log q^L(a^L) \approx \log q^0(a^0) - \epsilon \sum_{l=0}^{L-1} \text{Tr}(\nabla_{a^l} h(a^l)) + \mathcal{O}(\epsilon^2 dL). \quad (10)$$

Here, d is the dimensionality of a , *i.e.*, $a \in \mathbb{R}^d$ and $\mathcal{O}(\epsilon^2 dL)$ is the order of approximation error. *Proof Sketch:* As F is invertible, we apply the change of variable formula (Appendix C.2) on the transformation $F \circ F \circ \dots \circ F$ and obtain: $\log q^L(a^L) = \log q^0(a^0) - \sum_{l=0}^{L-1} \log |\det(I + \epsilon \nabla_{a^l} h(a^l))|$. Under the assumption $\epsilon \|\nabla_{a_i} h(a_i)\|_\infty \ll 1$, we apply the corollary of Jacobi’s formula (Appendix C.3) and get Eq. (10). The detailed proof is in Appendix F. Note that the condition $\epsilon \|\nabla_{a_i} h(a_i)\|_\infty \ll 1$ can always be satisfied when we choose a sufficiently small step size ϵ , or the gradient of $h(a)$ is small, *i.e.*, $h(a)$ is Lipschitz continuous with a sufficiently small constant. It follows from the theorem above, that the entropy of a policy modeled as an EBM sampler (Eq (9)) can be expressed analytically as:

$$\mathcal{H}(\pi_\theta(\cdot|s)) = -\mathbb{E}_{a_\theta^0 \sim q_\theta^0} \left[\log q_\theta^L(a_\theta^L|s) \right] \approx -\mathbb{E}_{a_\theta^0 \sim q_\theta^0} \left[\log q_\theta^0(a_\theta^0|s) - \epsilon \sum_{l=0}^{L-1} \text{Tr}(\nabla_{a_\theta^l} h(a_\theta^l, s)) \right]. \quad (11)$$

In the following, we drop the dependency of the action on θ for simplicity of the notation.

3.3 INVERTIBLE POLICIES

Next, we study the invertibility of three popular EBM samplers: SVGD, SGLD, and HMC as well as the efficiency of computing the trace, *i.e.*, $\text{Tr}(\nabla_{a^l} h(a^l, s))$ in Eq (10) for the ones that are invertible.

Proposition 3.2 (SVGD invertibility). *Given the SVGD learning rate ϵ and RBF kernel $k(\cdot, \cdot)$ with variance σ , if $\epsilon \ll \sigma$, the update rule of SVGD dynamics defined in Eq (1) is invertible.*

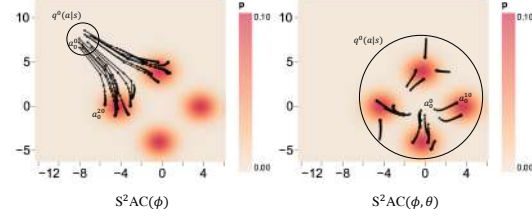


Figure 3: $S^2AC(\phi, \theta)$ achieves faster convergence to the target distribution (in orange) than $S^2AC(\phi)$ by parameterizing the initial distribution $\mathcal{N}(\mu_\theta, \sigma_\theta)$ of the SVGD sampler.

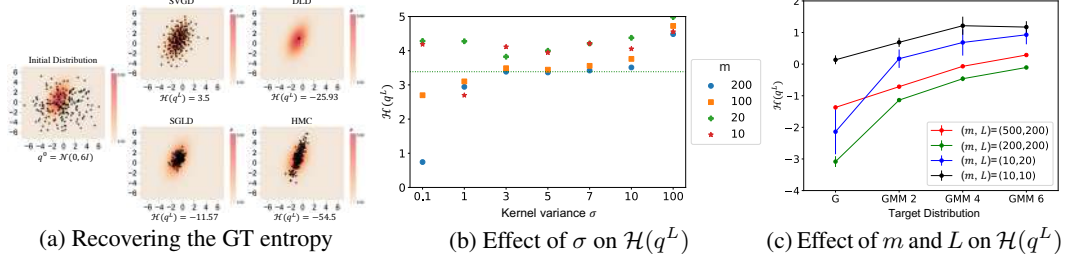


Figure 4: Entropy evaluation results.

Proof Sketch: We use the explicit function theorem to show that the Jacobian $\nabla_a F(a, s)$ of the update rule $F(a, s)$ is diagonally dominated and hence invertible. This yields invertibility of $F(a, s)$. See detailed proof in Appendix G.3.

Theorem 3.3. *The closed-form estimate of $\log q^L(a^L|s)$ for the SVGD based sampler with an RBF kernel $k(\cdot, \cdot)$ is*

$$\log q^L(a^L|s) \approx \log q^0(a^0|s) + \frac{\epsilon}{m\sigma^2} \sum_{l=0}^{L-1} \sum_{j=1, a^l \neq a_j^l}^m k(a_j^l, a^l) \left((a^l - a_j^l)^\top \nabla_{a_j^l} Q(s, a_j^l) + \frac{\alpha}{\sigma^2} \|a^l - a_j^l\|^2 - d\alpha \right).$$

Here, $(\cdot)^\top$ denotes the transpose of a matrix/vector. Note that the entropy does not depend on any matrix computation, but only on vector dot products and first-order vector derivatives. The proof is in Appendix H.1. Intuitively, the derived likelihood is proportional to (1) the concavity of the curvature of the Q-landscape, captured by a weighted average of the neighboring particles’ Q-value gradients and (2) pairwise-distances between the neighboring particles ($\sim \|a_i^l - a_j^l\|^2 \cdot \exp(\|a_i^l - a_j^l\|^2)$), i.e., the larger the distance the higher is the entropy. We elaborate on the connection between this formula and non-parametric entropy estimators in Appendix B.

Proposition 3.4 (SGLD, HMC). *The SGLD and HMC updates are not invertible w.r.t. a .*

Proof Sketch: SGLD is stochastic (noise term) and thus not injective. HMC is only invertible if conditioned on the velocity v . Detailed proofs are in Appendices G.1, G.2.

From the above theoretic analysis, we can see that SGLD update is not invertible and hence is not suitable as a sampler for S^2AC . While the HMC update is invertible, its derived closed-form entropy involves calculating Hessian and hence computationally more expensive. Due to these considerations, we choose to use SVGD with an RBF kernel as the underlying sampler of S^2AC .

4 RESULTS

We first evaluate the correctness of our proposed closed-form entropy formula. Then we present the results of different RL algorithms on multigoal and MuJoCo environments.

4.1 ENTROPY EVALUATION

This experiment tests the correctness of our entropy formula. We compare the estimated entropy for distributions (with known ground truth entropy or log-likelihoods) using different samplers and study the sensitivity of the formula to different samplers’ parameters. **(1) Recovering the ground truth entropy.** In Figure 4a, we plot samples (black dots) obtained by SVGD, SGLD, DLD and HMC at convergence to a Gaussian with ground truth entropy $\mathcal{H}(p) = 3.41$, starting from the same initial distribution (leftmost sub-figure). We also report the entropy values computed via Eq. (11). Unlike SGLD, DLD, and HMC, SVGD recovers the ground truth entropy. This empirically supports Proposition 3.4 that SGLD, DLD, and HMC are not invertible. **(2) Effect of the kernel variance.** Figure 4b shows the effect of different SVGD kernel variances σ , where we use the same initial Gaussian from Figure 4a. We also visualize the particle distributions after L SVGD steps for the different configurations in Figure 9 of Appendix I. We can see that when the kernel variance is too small (e.g., $\sigma = 0.1$), the invertibility is violated, and thus the estimated entropy is wrong even at convergence. On the other extreme when the kernel variance is too large (e.g., $\sigma = 100$), i.e., when the particles are too scattered initially, the particles do not converge to the target Gaussian due to noisy gradients in the first term of Eq. (11). The best configurations hence lie somewhere in between (e.g., $\sigma \in \{3, 5, 7\}$). **(3) Effect of SVGD steps and particles.** Figure 4c and Figure 10b (Appendix I) show the behavior of our entropy formula under different configurations of the number of SVGD steps and particles, on two settings: (i) GMM M with an increasing number of components M , and (ii) distributions with increasing ground truth entropy values, i.e., Gaussians with increasing variances σ . Results show that our entropy consistently grows with an increasing M (Figure 4c) and increasing σ (Figure 10b), even when a small number of SVGD steps and particles is used (e.g., $L = 10, m = 10$).

4.2 MULTI-GOAL EXPERIMENTS

To check if S^2AC learns a better solution to the max-entropy objective (Eq (2)), we design a new multi-goal environment as shown in Figure 5. The agent is a 2D point mass at the origin trying to reach one of the goals (in red). Q-landscapes are depicted by level curves. Actions are bounded in $[-1, 1]$ along both axes. Critical states for the analysis are marked with blue crosses. It is built on the multi-goal environment in Haarnoja et al. (2017) with modifications such that all the goals have (i) the same maximum expected future reward (positive) but (ii) different maximum expected future entropy. This is achieved by asymmetrically placing the goals (two goals on the left side and one on the right, leading to a higher expected future entropy on the left side) while assigning the same final rewards to all the goals. The problem setup and hyperparameters are detailed in Appendix J. **(1) Multi-modality.** Figure 6 visualizes trajectories (blue lines) collected from 20 episodes of $S^2AC(\phi, \theta)$, $S^2AC(\phi)$, SAC, SQL and SAC-NF (SAC with a normalizing flow policy, Mazouze et al. (2020)) agents (rows) at test time for increasing entropy weights α (columns). S^2AC and SQL consistently cover all the modes for all α values, while this is only achieved by SAC and SAC-NF for large α values. Note that, in the case of SAC, this comes at the expense of accuracy. Although normalizing flows are expressive enough in theory, they are known to quickly collapse to local optima in practice Kobyzev et al. (2020). The dispersion term in S^2AC encodes an inductive bias to mitigate this issue. **(2) Maximizing the expected future entropy.** We also see that with increasing α , more S^2AC and SAC-NF trajectories converge to the left goals (G_2/G_3). This shows both models learn to maximize the expected future entropy. This is not the case for SQL whose trajectory distribution remains uniform across the goals. SAC results do not show a consistent trend. This validates the hypothesis that the entropy term in SAC only helps exploration but does not lead to maximizing future entropy. The quantified distribution over reached goals is in Figure 12 of Appendix J. **(3) Robustness/adaptability.** To assess the robustness of the learned policies, we place an obstacle (red bar in Figure 7) on the path to G_2 . We show the test time trajectories of 20 episodes using S^2AC , SAC, SQL and SAC-NF agents trained with different α 's. We observe that, for S^2AC and SAC-NF, with increasing α , more trajectories reach the goal after hitting the obstacles. This is not the case for SAC, where many trajectories hit the obstacle without reaching the goal. SQL does not manage to escape the barrier even with higher α . Additional results on the **(4) effect of parameterization of q^0** , and the **(5) entropy's effect on the learned Q-landscapes** are respectively reported in Figure 11 and Figure 14 of Appendix J.

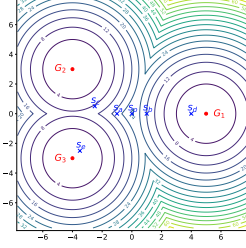


Figure 5: Multigoal env.

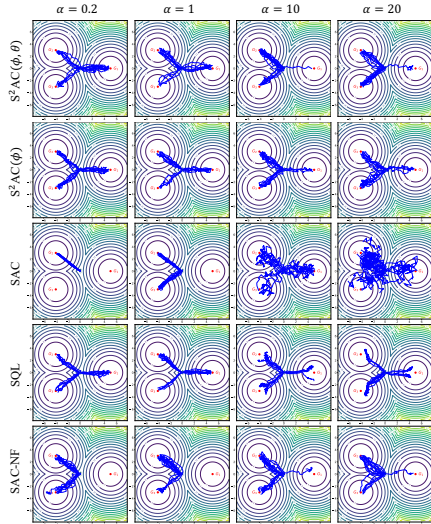


Figure 6: S^2AC and SAC-NF learn to maximize the expected future entropy (biased towards G_2/G_3) while SAC and SQL do not. S^2AC consistently recovers all modes, while SAC-NF with smaller α 's does not, indicating its instability.

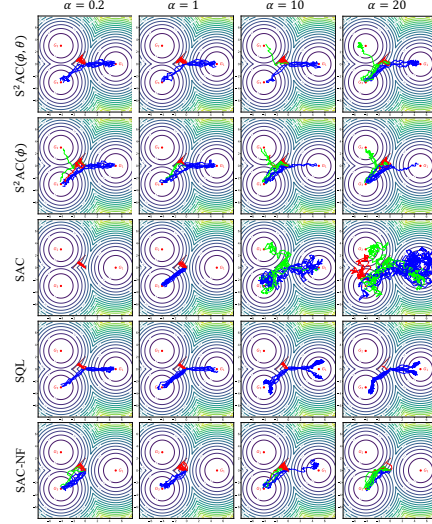


Figure 7: S^2AC and SAC-NF are more robust to perturbations. Obstacle O is placed diagonally at $[-1, 1]$. Trajectories that did and did not reach the goal after hitting O are in green and red, respectively.

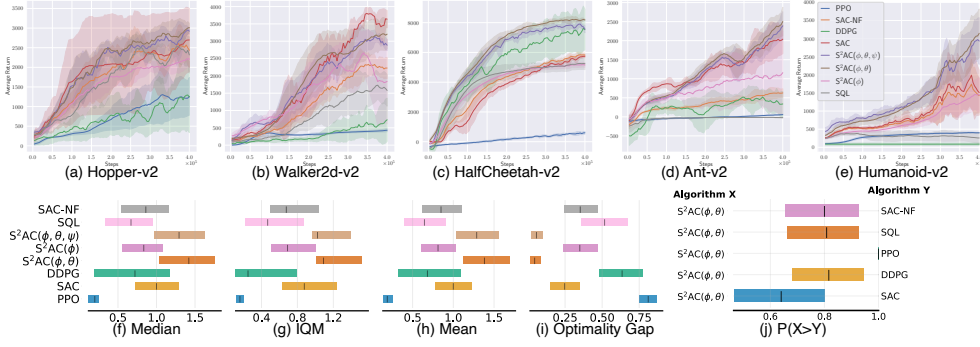


Figure 8: **(a)-(e)**: Performance curves on the MuJoCo benchmark (training). S^2AC outperforms SQL and SAC-NF on all environments and SAC on 4 out of 5 environments. **(f)-(i)**: Comparison of Median, IQM, Mean, and Optimality Gap between S^2AC and baseline algorithms. **(j)**: The probabilities of S^2AC outperforming baseline algorithms.

4.3 MUJOCO EXPERIMENTS

We evaluate S^2AC on five environments from MuJoCo (Brockman et al., 2016): Hopper-v2, Walker2d-v2, HalfCheetah-v2, Ant-v2, and Humanoid-v2. As baselines, we use (1) DDPG (Gu et al., 2017), (2) PPO (Schulman et al., 2015), (3) SQL (Haarnoja et al., 2017), (4) SAC-NF (Mazouze et al., 2020), and (5) SAC (Haarnoja et al., 2018a). Hyperparameters are in Appendix K.

(1) Performance and sample efficiency. We train five different instances of each algorithm with different random seeds, with each performing 100 evaluation rollouts every 1000 environment steps. Performance results are in Figure 8(a)-(e). The solid curves correspond to the mean returns over the five trials and the shaded region represents the minimum and maximum. $S^2AC(\phi, \theta)$ is consistently better than SQL and SAC-NF across all the environments and has superior performance than SAC in four out of five environments. Results also show that the initial parameterization was key to ensuring the scalability ($S^2AC(\phi)$ has poor performance compared to $S^2AC(\phi, \theta)$). Figure 8(f)-(j) demonstrate the statistical significance of these gains by leveraging statistics from the reliable library (Agarwal et al., 2021) which we detail in Appendix K.

(2) Run-time. We report the run-time of action selection of SAC, SQL, and S^2AC algorithms in Table 1. $S^2AC(\phi, \theta)$ run-time increases linearly with the action space. To improve the scalability, we train an amortized version that we deploy at test-time, following (Haarnoja et al., 2017). Specifically, we train a feed-forward deepnet $f_\psi(s, z)$ to mimic the SVGD dynamics during testing, where z is a random vector that allows mapping the same state to different particles. Note that we cannot use $f_\psi(s, z)$ during training as we need to estimate the entropy in Eq (11), which depends on the unrolled SVGD dynamics (details in Appendix K). The amortized version $S^2AC(\phi, \theta, \psi)$ has a similar run-time to SAC and SQL with a slight tradeoff in performance (Figure 8).

	Hopper	Walker2d	HalfCheetah	Ant
Action dim	3	6	6	8
State dim	11	17	17	111
SAC	0.723	0.714	0.731	0.708
SQL	0.839	0.828	0.815	0.836
$S^2AC(\phi, \theta)$	3.267	4.622	4.583	5.917
$S^2AC(\phi, \theta, \psi)$	0.850	0.817	0.830	0.837

Table 1: Action selection run-time on MuJoCo.

5 RELATED WORK

MaxEnt RL (Todorov, 2006; Ziebart, 2010; Rawlik et al., 2012) aims to learn a policy that gets high rewards while acting as randomly as possible. To achieve this, it maximizes the sum of expected future reward and expected future entropy. It is different from entropy regularization (Schulman et al., 2015; O’Donoghue et al., 2016; Schulman et al., 2017) which maximizes entropy at the current time step. It is also different from multi-modal RL approaches (Tang & Agrawal, 2018) which recover different modes with equal frequencies without considering their future entropy. MaxEnt RL has been broadly incorporated in various RL domains, including inverse RL (Ziebart et al., 2008; Finn et al., 2016), stochastic control (Rawlik et al., 2012; Toussaint, 2009), guided policy search (Levine & Koltun, 2013), and off-policy learning (Haarnoja et al., 2018a; b). MaxEnt RL is shown to maximize a lower bound of the robust RL objective (Eysenbach & Levine, 2022) and is hence less sensitive

to perturbations in state and reward functions. From the variational inference lens, MaxEnt RL aims to find the policy distribution that minimizes the KL -divergence to an EBM over Q-function. The desired family of variational distributions is (1) expressive enough to capture the intricacies of the Q-value landscape (e.g., multimodality) and (2) has a tractable entropy estimate. These two requirements are hard to satisfy. SAC (Haarnoja et al., 2018a) uses a Gaussian policy. Despite having a tractable entropy, it fails to capture arbitrary Q-value landscapes. SAC-GMM (Haarnoja, 2018) extends SAC by modeling the policy as a Gaussian Mixture Model, but it requires an impractical grid search over the number of components. Other extensions include IAPO (Marino et al., 2021) which also models the policy as a uni-modal Gaussian but learns a collection of parameter estimates (mean, variance) through different initializations. While this yields multi-modality, it does not optimize a MaxEnt objective. SSPG (Cetin & Celiktutan, 2022) and SAC-NF (Mazouze et al., 2020) respectively improve the policy expressivity by modeling the policy as a Markov chain with Gaussian transition probabilities and as a normalizing flow. Due to training instability, the reported multi-goal experiments in (Cetin & Celiktutan, 2022) show that, though both models capture multimodality, they fail to maximize the expected future entropy in positive reward setups. SQL (Haarnoja et al., 2017), on the other hand, bypasses the explicit entropy computation altogether via a soft version of value iteration. It then trains an amortized SVGD (Wang & Liu, 2016) sampler from the EBM over the learned Q-values. However, estimating soft value functions requires approximating integrals via importance sampling which is known to have high variance and poor scalability. We propose a new family of variational distributions induced by a parameterized SVGD sampler from the EBM over Q-values. Our policy is expressive and captures multi-modal distributions while being characterized by a tractable entropy estimate.

EBMs (LeCun et al., 2006; Wu et al., 2018) are represented as Gibbs densities $p(x) = \exp E(x)/Z$, where $E(x) \in \mathbb{R}$ is an energy function describing inter-variable dependencies and $Z = \int \exp E(x)$ is the partition function. Despite their expressiveness, EBMs are not tractable as the partition function requires integrating over an exponential number of configurations. Markov Chain Monte Carlo (MCMC) methods (Van Ravenzwaaij et al., 2018) (e.g., HMC (Hoffman & Gelman, 2014), SGLD (Welling & Teh, 2011)) are frequently used to approximate the partition function via sampling. There have been recent efforts to parameterize these samplers via deepnets (Levy et al., 2017; Gong et al., 2018; Feng et al., 2017) to improve scalability. Similarly to these methods, we propose a parameterized variant of SVGD (Liu & Wang, 2016) as an EBM sampler to enable scalability to high-dimensional action spaces. Beyond sampling, we derive a closed-form expression of the sampling distribution as an estimate of the EBM. This yields a tractable estimate of the entropy. This is opposed to previous methods for estimating EBM entropy which mostly rely on heuristic approximation, lower bounds (Dai et al., 2017; 2019a), or neural estimators of mutual information (Kumar et al., 2019). The idea of approximating the entropy of EBMs via MCMC sampling by leveraging the change of variable formula was first proposed in (Dai et al., 2019b). The authors apply the formula to HMC and LD, which, as we show previously, violate the invertibility assumption. To go around this, they augment the EBM family with the noise or velocity variable for LD and HMC respectively. But the derived log-likelihood of the sampling distribution turns out to be –counter-intuitively– independent of the sampler’s dynamics and equal to the initial distribution, which is then parameterized using a flow model (details in Appendix B.2). We show that SVGD is invertible, and hence we sample from the original EBM, so that our derived entropy is more intuitive as it depends on the SVGD dynamics. **SVGD-augmented RL** (Liu & Wang, 2016) has been explored under other RL contexts. Liu et al. (2017) use SVGD to learn a distribution over policy parameters. While this leads to learning diverse policies, it is fundamentally different from our approach as we are interested in learning a single multi-modal policy with a closed-form entropy formula. Castanet et al. (2023); Chen et al. (2021) use SVGD to sample from multimodal distributions over goals/tasks. We go beyond sampling and use SVGD to derive a closed-form entropy formula of an expressive variational distribution.

6 CONCLUSION

We propose S^2AC , an actor-critic algorithm that yields a more optimal solution to the MaxEnt RL objective than previously proposed approaches. S^2AC achieves this by leveraging a new family of variational distributions characterized by SVGD dynamics. The proposed distribution has high expressivity, i.e., it is flexible enough to capture multimodal policies in high dimensional spaces, and a tractable entropy estimate. Empirical results show that S^2AC learns expressive and robust policies while having superior performance than other MaxEnt RL algorithms. For future work, we plan to study the application of the proposed variational distribution to other domains and develop benchmarks to evaluate the robustness of RL agents.

ACKNOWLEDGMENTS

Bo An is supported by the National Research Foundation Singapore and DSO National Laboratories under the AI Singapore Programme (AISGAward No: AISG2-GC-2023-009). Haipeng Chen is supported by William & Mary FRC Faculty Research Grants.

REFERENCES

- Abbas Abdolmaleki, Jost Tobias Springenberg, Yuval Tassa, Remi Munos, Nicolas Heess, and Martin Riedmiller. Maximum a posteriori policy optimisation. *arXiv preprint arXiv:1806.06920*, 2018.
- Rishabh Agarwal, Max Schwarzer, Pablo Samuel Castro, Aaron Courville, and Marc G Bellemare. Deep reinforcement learning at the edge of the statistical precipice. *NeurIPS*, 2021.
- I. Ahmad and Pi-Erh Lin. A nonparametric estimation of the entropy for absolutely continuous distributions (corresp.). *IEEE Trans. Inf. Theory*, 1976.
- Jan Beirlant and M.C.A van Zuijlen. The empirical distribution function and strong laws for functions of order statistics of uniform spacings. *J. Multivar. Anal.*, 1985.
- Jan Beirlant, Edward J Dudewicz, László Györfi, Edward C Van der Meulen, et al. Nonparametric entropy estimation: An overview. *IJSRMSS*, 1997.
- Laura T Bernhofen, Edward J Dudewicz, Janos Levendovszky, and Edward C van der Meulen. Ranking of the best random number generators via entropy-uniformity theory. *AJMMS*, 1996.
- Peter J. Bickel and Leo Breiman. Sums of Functions of Nearest Neighbor Distances, Moment Bounds, Limit Theorems and a Goodness of Fit Test. *Ann. Probab.*, 1983.
- Greg Brockman, Vicki Cheung, Ludwig Pettersson, Jonas Schneider, John Schulman, Jie Tang, and Wojciech Zaremba. Openai Gym. *arXiv preprint arXiv:1606.01540*, 2016.
- Giuseppe Carleo and Matthias Troyer. Solving the quantum many-body problem with artificial neural networks. *Science*, 2017.
- Nicolas Castanet, Olivier Sigaud, et al. Stein variational goal generation for adaptive exploration in multi-goal reinforcement learning. 2023.
- Edoardo Cetin and Oya Celiktutan. Policy gradient with serial markov chain reasoning. *NeurIPS*, 2022.
- Jiayu Chen, Yuanxin Zhang, Yuanfan Xu, Huimin Ma, Huazhong Yang, Jiaming Song, Yu Wang, and Yi Wu. Variational automatic curriculum learning for sparse-reward cooperative multi-agent problems. *NeurIPS*, 2021.
- Thomas M Cover. *Elements of information theory*. John Wiley & Sons, 1999.
- Noel Cressie. Power results for tests based on high-order gaps. *Biometrika*, 1978.
- Bo Dai, Hanjun Dai, Arthur Gretton, Le Song, Dale Schuurmans, and Niao He. Kernel exponential family estimation via doubly dual embedding. In *AISTAT*, 2019a.
- Bo Dai, Zhen Liu, Hanjun Dai, Niao He, Arthur Gretton, Le Song, and Dale Schuurmans. Exponential family estimation via adversarial dynamics embedding. *NeurIPS*, 2019b.
- Zihang Dai, Amjad Almahairi, Philip Bachman, Eduard Hovy, and Aaron Courville. Calibrating energy-based generative adversarial networks. *arXiv preprint arXiv:1702.01691*, 2017.
- Jay L Devore, Kenneth N Berk, Matthew A Carlton, et al. *Modern mathematical statistics with applications*. Springer, 2012.
- Yan Duan, Xi Chen, Rein Houthooft, John Schulman, and Pieter Abbeel. Benchmarking deep reinforcement learning for continuous control. In *ICML*, 2016.
- Edward J Dudewicz and Edward C Van Der Meulen. Entropy-based tests of uniformity. *JASA*, 1981.

- Benjamin Eysenbach and Sergey Levine. Maximum entropy rl (provably) solves some robust rl problems. In *ICLR*, 2022.
- Yihao Feng, Dilin Wang, and Qiang Liu. Learning to draw samples with amortized stein variational gradient descent. *arXiv preprint arXiv:1707.06626*, 2017.
- Chelsea Finn, Sergey Levine, and Pieter Abbeel. Guided cost learning: Deep inverse optimal control via policy optimization. In *ICML*, 2016.
- Ruiqi Gao, Yang Song, Ben Poole, Ying Nian Wu, and Diederik P Kingma. Learning energy-based models by diffusion recovery likelihood. *arXiv preprint arXiv:2012.08125*, 2020.
- Xun Gao and Lu-Ming Duan. Efficient representation of quantum many-body states with deep neural networks. *Nature communications*, 2017.
- Wenbo Gong, Yingzhen Li, and José Miguel Hernández-Lobato. Meta-learning for stochastic gradient mcmc. *arXiv preprint arXiv:1806.04522*, 2018.
- Shixiang Gu, Timothy Lillicrap, Zoubin Ghahramani, Richard E Turner, and Sergey Levine. Q-prop: Sample-efficient policy gradient with an off-policy critic. In *ICLR*, 2017.
- László Györfi and Edward C. van der Meulen. Density-free convergence properties of various estimators of entropy. *CSDA*, 1987.
- Tuomas Haarnoja. *Acquiring diverse robot skills via maximum entropy deep reinforcement learning (Ph.D. thesis)*. University of California, Berkeley, 2018.
- Tuomas Haarnoja, Haoran Tang, Pieter Abbeel, and Sergey Levine. Reinforcement learning with deep energy-based policies. In *ICML*, 2017.
- Tuomas Haarnoja, Aurick Zhou, Pieter Abbeel, and Sergey Levine. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In *ICML*, 2018a.
- Tuomas Haarnoja, Aurick Zhou, Kristian Hartikainen, George Tucker, Sehoon Ha, Jie Tan, Vikash Kumar, Henry Zhu, Abhishek Gupta, Pieter Abbeel, et al. Soft actor-critic algorithms and applications. *arXiv preprint arXiv:1812.05905*, 2018b.
- Peter Hall. On powerful distributional tests based on sample spacings. *JMVA*, 1986.
- Hideitsu Hino and Noboru Murata. A conditional entropy minimization criterion for dimensionality reduction and multiple kernel learning. *Neural Comput.*, 2010.
- Matthew D Hoffman and Andrew Gelman. The No-U-Turn sampler: Adaptively setting path lengths in hamiltonian monte carlo. *JMLR*, 2014.
- Aleksandr Vasil’evich Ivanov and MN Rozhkova. On properties of the statistical estimate of the entropy of a random vector with a probability density. *Problemy Peredachi Informatsii*, 1981.
- Harry Joe. Estimation of entropy and other functionals of a multivariate density. *Ann. Inst. Stat. Math.*, 1989.
- Hilbert J Kappen. Path integrals and symmetry breaking for optimal control theory. *JSTAT*, 2005.
- Ivan Kobyzev, Simon JD Prince, and Marcus A Brubaker. Normalizing flows: An introduction and review of current methods. *PAMI*, 2020.
- Lyudmyla F Kozachenko and Nikolai N Leonenko. Sample estimate of the entropy of a random vector. *Problemy Peredachi Informatsii*, 1987.
- Rithesh Kumar, Anirudh Goyal, Aaron Courville, and Yoshua Bengio. Engan: Latent space mcmc and maximum entropy generators for energy-based models. *ICLR*, 2019.
- Alessandro Lazaric, Marcello Restelli, and Andrea Bonarini. Reinforcement learning in continuous action spaces through sequential monte carlo methods. *NeurIPS*, 2007.

- Erik G Learned-Miller and John W Fisher III. Ica using spacings estimates of entropy. *JMLR*, 2003.
- Yann LeCun, Sumit Chopra, Raia Hadsell, M Ranzato, and Fugie Huang. A tutorial on energy-based learning. *Predicting Structured Data*, 2006.
- Kyowoon Lee, Sol-A Kim, Jaesik Choi, and Seong-Wan Lee. Deep reinforcement learning in continuous action spaces: a case study in the game of simulated curling. In *ICML*, 2018.
- Sergey Levine and Vladlen Koltun. Guided policy search. In *ICML*, 2013.
- Daniel Levy, Matthew D Hoffman, and Jascha Sohl-Dickstein. Generalizing hamiltonian monte carlo with neural networks. *arXiv preprint arXiv:1711.09268*, 2017.
- Timothy P Lillicrap, Jonathan J Hunt, Alexander Pritzel, Nicolas Heess, Tom Erez, Yuval Tassa, David Silver, and Daan Wierstra. Continuous control with deep reinforcement learning. *arXiv preprint arXiv:1509.02971*, 2015.
- Haozhe Liu, Bing Li, Haoqian Wu, Hanbang Liang, Yawen Huang, Yuexiang Li, Bernard Ghanem, and Yefeng Zheng. Combating mode collapse in gans via manifold entropy estimation. *AAAI*, 2022.
- Qiang Liu. Stein variational gradient descent as gradient flow. *NeurIPS*, 2017.
- Qiang Liu and Dilin Wang. Stein variational gradient descent: a general purpose bayesian inference algorithm. In *NeurIPS*, 2016.
- Yang Liu, Prajit Ramachandran, Qiang Liu, and Jian Peng. Stein variational policy gradient. *UAI*, 2017.
- Jan R Magnus and Heinz Neudecker. *Matrix differential calculus with applications in statistics and econometrics*. John Wiley & Sons, 2019.
- Shie Mannor, Dori Peleg, and Reuven Rubinstein. The cross entropy method for classification. In *ICML*, 2005.
- Joseph Marino, Alexandre Piché, Alessandro Davide Ialongo, and Yisong Yue. Iterative amortized policy optimization. *NeurIPS*, 2021.
- Bogdan Mazouze, Thang Doan, Audrey Durand, Joelle Pineau, and R Devon Hjelm. Leveraging exploration in off-policy algorithms via normalizing flows. In *CoRL*, 2020.
- Roger G Melko, Giuseppe Carleo, Juan Carrasquilla, and J Ignacio Cirac. Restricted boltzmann machines in quantum physics. *Nature Physics*, 2019.
- Safa Messaoud. *Toward More Scalable Structured Models*. PhD thesis, University of Illinois Urbana-Champaign, 2021.
- Safa Messaoud, David Forsyth, and Alexander G Schwing. Structural consistency and controllability for diverse colorization. In *ECCV*, pp. 596–612, 2018.
- Safa Messaoud, Maghav Kumar, and Alexander G Schwing. Can we learn heuristics for graphical model inference using reinforcement learning? In *CVPR Workshops*, pp. 766–767, 2020.
- Volodymyr Mnih, Koray Kavukcuoglu, David Silver, Andrei A Rusu, Joel Veness, Marc G Bellemare, Alex Graves, Martin Riedmiller, Andreas K Fidjeland, Georg Ostrovski, et al. Human-level control through deep reinforcement learning. *Nature*, 2015.
- Radford M Neal et al. Mcmc using hamiltonian dynamics. *Handbook of Markov Chain Monte Carlo*, 2011.
- Brendan O’Donoghue, Remi Munos, Koray Kavukcuoglu, and Volodymyr Mnih. Combining policy gradient and q-learning. *arXiv preprint arXiv:1611.01626*, 2016.
- Bo Pang, Tianyang Zhao, Xu Xie, and Ying Nian Wu. Trajectory prediction with latent belief energy-based model. In *CVPR*, pp. 11814–11824, 2021.

- Liam Paninski. Estimation of entropy and mutual information. *Neural Comput.*, 2003.
- Fernando Pérez-Cruz. Estimation of information theoretic measures for continuous random variables. In *NeurIPS*, 2008.
- Konrad Rawlik, Marc Toussaint, and Sethu Vijayakumar. On stochastic optimal control and reinforcement learning by approximate inference. *Proceedings of Robotics: Science and Systems VIII*, 2012.
- Danilo Rezende and Shakir Mohamed. Variational inference with normalizing flows. In *ICML*. PMLR, 2015.
- Édgar Roldán, Jérémie Barral, Pascal Martin, Juan MR Parrondo, and Frank Jülicher. Quantifying entropy production in active fluctuations of the hair-cell bundle from time irreversibility and uncertainty relations. *New J. Phys.*, 2021.
- Reuven Y. Rubinstein and Dirk P. Kroese. *The Cross Entropy Method: A Unified Approach To Combinatorial Optimization, Monte-Carlo Simulation (Information Science and Statistics)*. Springer-Verlag, 2004.
- Ruslan Salakhutdinov, Andriy Mnih, and Geoffrey Hinton. Restricted boltzmann machines for collaborative filtering. In *ICML*, 2007.
- John Schulman, Sergey Levine, Pieter Abbeel, Michael Jordan, and Philipp Moritz. Trust region policy optimization. In *ICML*, 2015.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Claude Elwood Shannon. A mathematical theory of communication. *ACM SIGMOBILE*, 2001.
- David Silver, Guy Lever, Nicolas Heess, Thomas Degris, Daan Wierstra, and Martin Riedmiller. Deterministic policy gradient algorithms. In *ICML*, 2014.
- Richard S Sutton, David McAllester, Satinder Singh, and Yishay Mansour. Policy gradient methods for reinforcement learning with function approximation. *NeurIPS*, 1999.
- Yunhao Tang and Shipra Agrawal. Boosting trust region policy optimization by normalizing flows policy. *arXiv preprint arXiv:1809.10326*, 2018.
- F.P. Tarasenko. On the evaluation of an unknown probability density function, the direct estimation of the entropy from independent observations of a continuous random variable, and the distribution-free entropy test of goodness-of-fit. *Proceedings of the IEEE*, 1968.
- Evangelos Theodorou, Jonas Buchli, and Stefan Schaal. A generalized path integral control approach to reinforcement learning. *JMLR*, 2010.
- Emanuel Todorov. Linearly-solvable markov decision problems. In *NeurIPS*, 2006.
- Giacomo Torlai, Guglielmo Mazzola, Juan Carrasquilla, Matthias Troyer, Roger Melko, and Giuseppe Carleo. Neural-network quantum state tomography. *Nature Physics*, 2018.
- Marc Toussaint. Robot trajectory optimization using approximate inference. In *ICML*, 2009.
- Don Van Ravenzwaaij, Pete Cassey, and Scott D Brown. A simple introduction to markov chain monte-carlo sampling. *Psychon. Bull. Rev.*, 2018.
- Oldrich Vasicek. *A Test for Normality Based on Sample Entropy*. JSTOR, 2015.
- Nino Vieillard, Olivier Pietquin, and Matthieu Geist. Munchausen reinforcement learning. In *NeurIPS*, 2020.
- Dilin Wang and Qiang Liu. Learning to draw samples: With application to amortized mle for generative adversarial learning. *arXiv preprint arXiv:1611.01722*, 2016.

- Max Welling and Yee W Teh. Bayesian learning via stochastic gradient langevin dynamics. In *ICML*, 2011.
- Ying Nian Wu, Jianwen Xie, Yang Lu, and Song-Chun Zhu. Sparse and deep generalizations of the frame model. *Annals of Mathematical Sciences and Applications*, 2018.
- Markus Wulfmeier, Peter Ondruska, and Ingmar Posner. Maximum entropy deep inverse reinforcement learning. *arXiv preprint arXiv:1507.04888*, 2015.
- Jianwen Xie, Yang Lu, Ruiqi Gao, Song-Chun Zhu, and Ying Nian Wu. Cooperative training of descriptor and generator networks. *IEEE PAMI*.
- Jianwen Xie, Yang Lu, Song-Chun Zhu, and Yingnian Wu. A theory of generative convnet. In *ICML*, 2016.
- Jianwen Xie, Zilong Zheng, Ruiqi Gao, Wenguan Wang, Song-Chun Zhu, and Ying Nian Wu. Generative voxelnet: learning energy-based models for 3d shape synthesis and analysis. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(5), 2020.
- Jianwen Xie, Zilong Zheng, Xiaolin Fang, Song-Chun Zhu, and Ying Nian Wu. Learning cycle-consistent cooperative networks via alternating mcmc teaching for unsupervised cross-domain translation. In *AAAI*, volume 35, pp. 10430–10440, 2021a.
- Jianwen Xie, Zilong Zheng, and Ping Li. Learning energy-based model with variational auto-encoder as amortized sampler. In *AAAI*, 2021b.
- Jianwen Xie, Yaxuan Zhu, Jun Li, and Ping Li. A tale of two flows: Cooperative learning of langevin flow and normalizing flow toward energy-based model. *arXiv preprint arXiv:2205.06924*, 2022.
- Yang Zhao, Jianwen Xie, and Ping Li. Learning energy-based generative models via coarse-to-fine expanding and sampling. In *ICLR*, 2021.
- Zilong Zheng, Jianwen Xie, and Ping Li. Patchwise generative convnet: Training energy-based models from a single natural image for internal learning. In *CVPR*, pp. 2961–2970, June 2021.
- Mo Zhou and Jianfeng Lu. Single timescale actor-critic method to solve the linear quadratic regulator with convergence guarantees. *JMLR*, 2023.
- Brian D Ziebart. *Modeling purposeful adaptive behavior with the principle of maximum causal entropy*. Carnegie Mellon University, 2010.
- Brian D Ziebart, Andrew L Maas, J Andrew Bagnell, Anind K Dey, et al. Maximum entropy inverse reinforcement learning. In *AAAI*, 2008.