

Meta Variational Quantum Monte Carlo

Tianchen Zhao¹, James Stokes², Shravan Veerapaneni^{1,2,*}

¹Department of Mathematics, University of Michigan, Ann Arbor, MI 48109

²Flatiron Institute, Simons Foundation, New York, NY 10010

*Corresponding author. *Email:* shravan@umich.edu

Abstract

Motivated by close analogies between meta reinforcement learning (Meta-RL) and variational quantum Monte Carlo with disorder, we propose a learning problem and an associated notion of generalization, with applications in ground state determination for quantum systems described by random Hamiltonians. Specifically, we elaborate on a proposal of Zhao *et al.* (2020) interpreting the Hamiltonian disorder as task uncertainty for a Meta-RL agent. A model-agnostic meta-learning approach is proposed to solve the associated learning problem and numerical experiments in disordered quantum spin systems indicate that the resulting Meta Variational Monte Carlo accelerates training and [improves converged energies](#).

Keywords: Meta-learning, Variational Monte Carlo, Neural Quantum States

1 Introduction

Although deep neural networks excel in individual learning tasks, they are brittle with respect to task deformation. This fragility presents a challenge to the design of artificially intelligent agents which are required to efficiently adapt from known source tasks to a stream of unknown and dynamically changing target tasks. In order to quantify the ability of an agent to adapt when confronted with a stream of learning tasks, it proves convenient to adopt the modeling assumption in which there exists a probability distribution supported on the space of possible tasks called the task distribution. Meta-learning (also called *learning to learn* [1]) attempts to formalize the goal of adaptivity by exploiting regularities in the task distribution in order to output a hypothesis

that performs well on new task realizations, given limited data access to each target task. It is instructive to contrast meta-learning with the comparatively simpler paradigm of transfer learning in which the target task is *known*. In this paper, we explore a formal identification between model-agnostic meta-learning [2] and a seemingly unrelated problem in quantum physics involving quantum Monte Carlo with disorder.

Given a fixed target Hamiltonian H acting on a Hilbert space $\mathcal{H}$ of exponentially high dimension, Variational Monte Carlo (VMC) [3] computes an estimate of its minimal eigenvalue and a description of the associated eigenvector. The ability of VMC to overcome the curse-of-dimensionality in this context stems from a reformulation of the Rayleigh-Ritz principle as a stochastic optimization problem, the optimum of which is a function outputting the components of the ground eigenvector in some orthonormal basis. Leveraging the close connection between VMC and deep reinforcement learning, Carleo and Troyer [4] showed that when neural networks are exploited as trial functions and optimized using natural gradient techniques, VMC can achieve state-of-the-art results in finding the ground state energies of the antiferromagnetic Heisenberg model, in which the Hamiltonian is defined by a geometrically local interaction graph corresponding to a two-dimensional lattice. The domain of applicability of so-called neural-network quantum states has since been expanded to encompass problems of electronic structure in finite [5–7] and infinite dimensions [8]. Further connections between VMC and deep learning have been elaborated in [9, 10] where it was shown that VMC is a quantum generalization of Natural Evolution Strategies (NES) [11], which in turn, provides a single-step realization of natural policy gradient learning [12].

The ability of VMC to obtain state-of-the-art results comes at the expense of significant computation time due to the sequential nature of the sampling process. Sharir *et al.* [13] proposes autoregressive modelling in replacement of the original Markov chain Monte Carlo (MCMC) to speed-up the sampling, and Zhao *et al.* [14] focuses on the scalability of VMC to large problems and efficient use of all available resources. [This work](#) attempts to address the similar problem but under a different context, where we assume the presence of a sparse matrix ensembles that admit a certain task regularity. Our hypothesis is that one can accelerate the training of VMC and [improve the converged energies](#) on new learning tasks, by employing information from previously encountered tasks. The naive approach that uses the pretrained model parameters from one task as the initialization on the target task fails to apply in VMC, as there is no notion of generalization to out-of-sample data for a task of which the stochastic objective function is an unbiased estimator of a given population objective. This lack of delineation between training and testing phase is closely analogous to deep reinforcement learning, where agents are trained and tested in the same learning environment and the algorithm typically outputs a policy that is strongly overfitted to the learning task. The above considerations motivate the viewpoint that meta-learning provides the relevant context in which to discuss generalization both for deep reinforcement

learning agents and VMC. Indeed, Meta-RL has enjoyed significant progress in the last few years, propelled by the discovery of a scalable gradient-based instantiation suitable for deep learning called model-agnostic meta-learning [2] (MAML).

In this work, we investigate the empirical performance of meta-VMC over sparse matrix ensembles with different kinds of task regularity, which we encode via geometric locality assumptions. Our experimental results suggest that meta-VMC is capable of effectively accelerating the training of VMC and [improving the converged energies on new task realizations](#) from the given ensemble. This paper expands on the preliminary work of Zhao *et al.* [15] in which the notion of meta-VMC was introduced and numerically investigated for diagonal matrix ensembles as a special case, which correspond to the Max-Cut optimization problem. The content of the paper is organized as follows: in section 2, we introduce single-task variational Monte Carlo, emphasizing the connection with the REINFORCE [16, 17] algorithm and natural policy gradient learning. The basics of meta-learning and the meta-VMC are then recalled. The theory underlying model-agnostic meta-learning and gradient-based meta-VMC are described in section 3. Experimental results and analysis are presented in section 4 and section 5 concludes the paper.

2 Background

2.1 Variational Monte Carlo

This paper proposes heuristic approximation algorithms for determining a minimal eigenpair of certain large and sparse random conjugate-symmetric (Hermitian) matrices that admit an efficient description. For simplicity, we restrict to the case of real symmetric matrices and consider only real eigenvectors. In addition, the ground eigenvector is non-negative entry-wise if all off-diagonal entries of H are further restricted to be non-positive, as a consequence of the Perron-Frobenius theorem. The $N \times N$ matrices under consideration are required to be s -sparse in the sense the number of nonzero entries per row is $O(s)$ for $s \ll N$. In this work, we consider s -sparse quantum many-body Hamiltonians where $s = O(\text{poly}(\log(N)))$, although the techniques we discuss do not require an exponential separation between the sparsity parameter s and the matrix side length N . In addition, we require that for each row $x \in [N]$ of the matrix H , the list of nonzero entries and their locations $\{(y, H_{xy}) : H_{xy} \neq 0\}$ can be determined in time $O(s)$.

We describe the differentiable family of trial vectors $\theta \in \mathbb{R}^d$ with a function $\psi_\theta : [N] \rightarrow \mathbb{C}$, of which the outputs are the components of the vector relative to the standard basis $\psi_\theta(x) = \langle e_x, \psi_\theta \rangle$. Given a row-sparse Hermitian matrix H , we define the variational Monte Carlo learning problem as the following continuous stochastic optimization task,

$$\min_{\theta \in \mathbb{R}^d} L(\theta) \ , \quad L(\theta) := \frac{\langle \psi_\theta, H \psi_\theta \rangle}{\langle \psi_\theta, \psi_\theta \rangle} = \mathbb{E}_{x \sim \pi_\theta} \left[\frac{(H \psi_\theta)(x)}{\psi_\theta(x)} \right] \geq \lambda_{\min}(H) \ , \quad (1)$$

where the population quantity is computed over the probability distribution

$$\pi_\theta(x) := \frac{|\psi_\theta(x)|^2}{\langle \psi_\theta, \psi_\theta \rangle} . \quad (2)$$

It is convenient to express the population objective function as the mean of a random variable as $L(\theta) = \mathbb{E}_{x \sim \pi_\theta} [l_\theta(x)]$ where we have defined the following stochastic objective function, which is defined for $x \in [N]$ whenever $\psi_\theta(x) \neq 0$,

$$l_\theta(x) := \frac{(H\psi_\theta)(x)}{\psi_\theta(x)} , \quad (3)$$

and whose variance under π_θ is given by,

$$\text{var}_{x \sim \pi_\theta} (l_\theta(x)) := \mathbb{E}_{x \sim \pi_\theta} [|l_\theta(x) - L(\theta)|^2] = \frac{\langle \psi_\theta, H^2 \psi_\theta \rangle}{\langle \psi_\theta, \psi_\theta \rangle} - \left[\frac{\langle \psi_\theta, H \psi_\theta \rangle}{\langle \psi_\theta, \psi_\theta \rangle} \right]^2 . \quad (4)$$

It follows from the Rayleigh-Ritz principle that if the trial vector ψ_θ approaches any eigenvector of H , then the variance of stochastic objective approaches zero. In practice, the objective function is optimized using a variant of stochastic mini-batch gradient descent closely related to the stochastic natural gradient called stochastic reconfiguration [18]. Stochastic estimators for the gradient and the Fisher information matrix follow from their population forms. In the special case of real-valued wavefunctions,

$$\begin{aligned} \nabla L(\theta) &= 2 \mathbb{E}_{x \sim \pi_\theta} [(l_\theta(x) - L(\theta)) \nabla_\theta \log |\psi_\theta(x)|] , \\ I(\theta) &= \mathbb{E}_{x \sim \pi_\theta} [\nabla_\theta \log \pi_\theta(x) \otimes \nabla_\theta \log \pi_\theta(x)] , \end{aligned} \quad (5)$$

where we used an identity for the derivative of the logarithm. If the normalizing constant $\langle \psi_\theta, \psi_\theta \rangle$ of the probability distribution π_θ is unknown, then above expectation values can be approximated by the Markov chain Monte Carlo method. If the normalization condition $\langle \psi_\theta, \psi_\theta \rangle = 1$ is fulfilled, then by absorbing the factor of 2 into the logarithm one finds,

$$\nabla L(\theta) = \mathbb{E}_{x \sim \pi_\theta} [(l_\theta(x) - L(\theta)) \nabla_\theta \log \pi_\theta(x)] . \quad (6)$$

In the special case where H is diagonal, it was noted in [10] that $l_\theta(x)$ becomes independent of θ and thus the VMC algorithm can be understood either as natural evolution strategies [10] or equivalently as single-step natural policy gradient, also known as REINFORCE [16, 17].

2.2 Meta-learning

In contrast to the single-task formulation of VMC, which accepts a fixed target Hamiltonian as input, our proposed *meta-VMC* asks for an approximation of

the ground energy for an *ensemble* of Hamiltonians H_τ indexed by a random disorder parameter τ sampled from a known distribution $\mathcal{T}$. The simplest strategy of retraining a separate neural-network quantum state from scratch for each realization of the disorder parameter τ is impractical. The goal is thus shifted to finding a neural network that is maximally adaptive to new realizations of the disorder. For experimental support, our focus in this paper is a special case of disordered quantum spin systems. These results are viewed as a stepping stone to random electronic structures, in which we anticipate similar optimization considerations to apply.

The formulation of meta-VMC exhibits obvious parallels with *meta-learning* or *learning-to-learn* in the machine learning literature [1], where data from previously encountered learning tasks is employed to accelerate performance on new tasks, drawn from an underlying task distribution $\mathcal{T}$. In the language of meta-learning, τ indexes the learning task and $\mathcal{T}$ denotes the distribution over all tasks. In meta-VMC, we assume the task distribution is known to the learner; in contrast, conventional meta-learning assumes $\mathcal{T}$ is unknown but possesses sufficient regularity to render meta-learning feasible.

2.3 Relationship with previous work

In this section, we differentiate our proposal from the uses of meta-learning that have been proposed elsewhere in the quantum information literature. In [19], for example, meta-learning has been proposed to mitigate various sources of noise, specifically shot noise and parameter noise. In the context of VMC, shot noise is analogous to variance associated with finite mini-batches, whereas parameter noise has no clear analogue. Ref. [20] is the most similar to ours in that they consider meta-learning from known distributions. They differ by the choice to focus on variational quantum algorithms such as VQE and QAOA and by the fact that they do not use model-agnostic meta-learning. Instead, the meta-learning outer-loop involves training a separate recurrent neural network, similar to [21]. [In our work, we do not take this learning-to-learn approach of training an adaptive optimizer, but rather learn an initialization that can accelerate subsequent standard training on different but related tasks. The former approach is more suitable for example in large-scale computer vision tasks where the size of the LSTM optimizer is negligible in comparison with that of a CNN network; such an optimizer would be likely lead to severe overfitting for our tasks of interest.](#) The notion of meta-VMC and a model-agnostic training algorithm was originally introduced in [15]. The numerical experiments of [15] restricted to diagonal matrix ensembles which can be understood as classical combinatorial optimization problems. In this paper, we expand the experiments to include off-diagonal matrix ensembles which have no classical analog, since their ground eigenfunction involves a superposition of basis states.

3 Theory

A simple strategy that has proven successful in meta-learning of deep neural networks is multi-task transfer learning [22, 23], which aims to learn an initialization for subsequent tasks by jointly optimizing the learning objective of multiple tasks simultaneously, using a mini-batch training strategy that interleaves batches across the tasks. Multi-task learning is, however, prone to catastrophic interference [24], making it unsuitable for generalization to the VMC. The problem is exemplified by some of the simplest examples of disordered spin systems: suppose H_τ is a random Hamiltonian whose expected value under the disorder parameter vanishes $\mathbb{E}_{\tau \sim \mathcal{T}}[H_\tau] = 0$. As a concrete example, consider the Sherrington-Kirkpatrick Hamiltonian, in which τ represents a collection of i.i.d. centered Gaussian random variables $J_{ij} \sim N(0, 1)$ representing the exchange energies. If we denote by L_τ the objective function corresponding to disorder parameter τ , then the multi-task learning objective function, expressed in the population limit, is given by

$$L_{\text{MTL}}(\theta) := \mathbb{E}_{\tau \sim \mathcal{T}} [L_\tau(\theta)] = \mathbb{E}_{\tau \sim \mathcal{T}} \left\{ \mathbb{E}_{x \sim \pi_\theta} \left[\frac{(H_\tau \psi_\theta)(x)}{\psi_\theta(x)} \right] \right\} = \frac{\langle \psi_\theta, \mathbb{E}[H_\tau] \psi_\theta \rangle}{\langle \psi_\theta, \psi_\theta \rangle} = 0 . \quad (7)$$

The fact that the multi-task learning objective loses dependence on θ in the population limit implies that the associated mini-batch algorithm makes no progress asymptotically.

In order to define an objective function which is asymptotically non-vacuous and which promotes adaptation to new realizations of disorder, we propose to optimize the following meta-learning objective function, again presented in population form for simplicity [2, 21],

$$L_{\text{ML}}(\theta) := \mathbb{E}_{\tau \sim \mathcal{T}} [L_\tau(U_\tau^t(\theta))] = \mathbb{E}_{\tau \sim \mathcal{T}} [L_\tau(\underbrace{U_\tau \circ \dots \circ U_\tau}_{t \text{ times}}(\theta))] , \quad (8)$$

where $U_\tau^t : \mathbb{R}^d \rightarrow \mathbb{R}^d$ denotes the t -fold application of a task adaptation operator U_τ , which in the simplest case of gradient descent with step size β , is given by $U_\tau(\theta) = \theta - \beta \nabla L_\tau(\theta)$. Optimization of the meta-learning objective L_{ML} ensures that when a new realization of the disorder parameter is drawn, the initialization performs well after performing one or more steps of gradient descent. Loosely speaking, meta-learning can be justified when one has a budget for running a few steps of gradient descent. **In practice, we train a model using a finite number of training tasks generated from a given task distribution, and use the resulting parameters to initialize subsequent training on new test tasks drawn from the same distribution.**

In the case of meta-VMC, we consider gradient-based optimization. Specifically, we focus on model-agnostic meta-learning (MAML) [2] which is a

Input: Matrix ensemble $\mathcal{T}$, adaptation operator U_τ , adaptation steps t
Initialize θ ;
while *not done* **do**
 Sample batch of disorder parameters $B \stackrel{\text{iid}}{\sim} \mathcal{T}$;
 for *each disorder parameter* $\tau \in B$ **do**
 $\theta_\tau = U_\tau^t(\theta)$;
 $\nabla_\tau = (U_\tau^t)'(\theta) \nabla L_\tau(\theta_\tau)$;
 end
 $\nabla = \frac{1}{|B|} \sum_{\tau \in B} \nabla_\tau$;
 $\theta \leftarrow \text{OPTIMIZER}(\theta, \nabla)$;
end

Algorithm 1: MAML [2] adapted to meta-VMC (batched over tasks).

gradient-based algorithm that has been proposed for optimizing the meta-learning objective. Straightforward application of the chain rule gives rise to the following gradient estimator for the meta-learning objective,

$$\nabla L_{\text{ML}}(\theta) = \mathbb{E}_{\tau \sim \mathcal{T}} [(U_\tau^t)'(\theta) \nabla L_\tau(U_\tau^t(\theta))] , \quad (9)$$

where $(U_\tau^t)'(\theta)$ denotes the Jacobian matrix of the function $U_\tau^t : \mathbb{R}^d \rightarrow \mathbb{R}^d$. The pseudocode for MAML is outlined in Algorithm 1. In order to facilitate readability, we have presented the algorithm with batching only in the task index, leaving the remaining expectation values (with respect to Born probabilities) in population form. In a practical algorithm, the intermediate variables θ_τ and ∇_τ are estimated stochastically using independent batches of data generated by the same task τ . Since the computation of the Jacobian involves an expensive back-propagation, first-order MAML (foMAML) has been proposed (e.g., [2, 25]) as a simplification of MAML, in which the Jacobian matrix is approximated by the identity matrix.

3.1 An illustrative example

The fact that the meta-learning objective function manages to avoid the catastrophic interference phenomenon can be illustrated by the following toy model¹. Rather than considering the Rayleigh quotient, consider the following ensemble of quadratic functions specified by a random positive-definite matrix $A \in \mathbb{R}^{d \times d}$ and a random vector $b \in \mathbb{R}^d$,

$$L_\tau(\theta) = \frac{1}{2} \langle \theta, A \theta \rangle - \langle b, \theta \rangle , \quad (10)$$

where the random variable $\tau = (A, b)$ now corresponds to the task label. In the simplest setting of single-step ($t = 1$) meta-learning with vanilla update

¹This quadratic model has also been analyzed in the context of convergence theory in [26].

operator $U_\tau(\theta) = \theta - \beta \nabla L_\tau(\theta)$, the optimal solution of the multi-task and meta-learning objectives can be found in closed form,

$$\arg \min_{\theta \in \mathbb{R}^d} L_{\text{ML}}(\theta) = \mathbb{E}[A(I - \beta A)^2]^{-1} \mathbb{E}[(I - \beta A)^2 b] . \quad (11)$$

In the limit $\beta \rightarrow 0$ corresponding to multi-task learning, the optimal solution is found to only depend on the mean value of the random variable τ , whereas the meta-learner (corresponding to $\beta > 0$) exploits information in the higher-order moments of τ .

4 Experiments

For our experiments, we focus on symmetric matrices whose side length is a power of 2, that is, $N = 2^n$. The real vector space of $2^n \times 2^n$ symmetric matrices contains a subspace of dimension $n(n+3)/2 = O(\text{poly}(n))$ which is parametrized by real parameters $g_i, h_i, g_{ij} \in \mathbb{R}$ as follows,

$$H = - \sum_{1 \leq i \leq n} (h_i X_i + g_i Z_i) - \sum_{1 \leq i < j \leq n} g_{ij} Z_i Z_j , \quad (12)$$

where $X_i := I^{\otimes(i-1)} \otimes X \otimes I^{\otimes(n-i)}$ and $Z_i := I^{\otimes(i-1)} \otimes Z \otimes I^{\otimes(n-i)}$ are defined in terms of the following elementary 2×2 matrices,

$$I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} , \quad Z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} , \quad X = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} . \quad (13)$$

The Hamiltonian H can be verified to be n -sparse, and admits a binary representation with entry value written as

$$\begin{aligned} H_{xy} = & - \sum_{1 \leq i \leq n} h_i \delta_{x_1 y_1} \cdots \delta_{\neg x_i y_i} \cdots \delta_{x_n y_n} - \delta_{xy} \sum_{1 \leq i \leq n} g_i (1 - 2x_i) \\ & - \delta_{xy} \sum_{1 \leq i < j \leq n} g_{ij} (1 - 2x_i)(1 - 2x_j) , \end{aligned} \quad (14)$$

where the row and column indices x, y are in their binary forms $x = 2^{n-1}x_1 + \cdots + 2^0x_n, y = 2^{n-1}y_1 + \cdots + 2^0y_n$, and $\neg x_i$ denotes logical negation of $x_i \in \{0, 1\}$.

Given a description of the matrix ensemble and a differentiable family of trial vectors described by a neural network, we seek an initialization strategy which rapidly accelerates the convergence of the trial vector to the ground space of randomly drawn problem instances. The proposed strategy was compared against training random problem instances from scratch (randomly initialized neural network). We also consider random initialized models that are trained using the stochastic reconfiguration method (stochastic natural

gradient descent), as an additional baseline. To showcase the robustness of meta-VMC, we conduct experiments across various settings involving different task distributions and model architectures.

4.1 Task distribution generation

Matrix ensembles of exponential size were chosen by specifying the distribution for the $O(\text{poly}(n))$ parameters $g_i, h_i, g_{i,j}$ appearing in (14). In this work, we impose $h_i \leq 0$ to ensure that the ground eigenvector can be chosen to be a non-negative vector due to the Perron-Frobenius theorem. Meta-learning experiments were conducted using four task distributions (matrix ensembles) supported on the subspace of symmetric matrices of the form (14), by first fixing a base matrix in and then introducing Gaussian noise to the parameters. The first base matrix we consider is referred to as the Max-Cut problem and is defined by

$$H_{\text{Max-Cut}} = -\frac{1}{4} \sum_{1 \leq i < j \leq n} L_{ij} Z_i Z_j, \quad (15)$$

where $L = [L_{ij}]$ denotes the Laplacian matrix for an undirected graph $G = (V, E)$ of size $V = n$. The diagonal entries of this matrix correspond to the sizes of the 2^n possible cuts on the graph G . The adjacency matrix for G was chosen by forming the $n \times n$ matrix $(B + B^T)/2 - \text{diag}(B)$ with entries B_{ij} s which are randomly generated bits. By design, the diagonal entries of G are all zero.

The second base matrix is referred to as Sherrington-Kirkpatrick model, which is a matrix of the form (14) with $g_i, g_{ij} \sim U(-1, 1)$ and $h_i \sim U(0, 1)$ sampled once and fixed. The final experiments consider geometrically local versions of the Sherrington-Kirkpatrick model (referred to as transverse field Ising model), in which the interactions are determined by a one-dimensional ring geometry $\mathbb{Z}_L = \{0, \dots, L-1\}$ (with addition defined modulo L) and a two-dimensional torus geometry $\mathbb{Z}_L \times \mathbb{Z}_L$. The respective base matrices are given by

$$H_{\text{TIM-1D}} = - \sum_{i \in \mathbb{Z}_L} (g_i Z_i Z_{i+1} + h_i X_i), \quad (16)$$

$$H_{\text{TIM-2D}} = - \sum_{(i,j) \in \mathbb{Z}_L^2} (g_{i,j}^v Z_{i,j} Z_{i+1,j} + g_{i,j}^h Z_{i,j} Z_{i,j+1} + h_{i,j} X_{i,j}). \quad (17)$$

Having fixed the base matrices as above, sampling from the respective matrix ensembles was achieved via the following procedure. In the case of Max-Cut, a random adjacency matrix was formed by perturbing the base adjacency matrix A with additive noise matrix δA and then rebinarizing the sum $A + \delta A$. In particular, we chose $\delta A = (N + N^T)/2$ where the entries of N consist of independent centered Gaussian noise of variance σ^2 . In the case of the transverse field Ising and Sherrington-Kirkpatrick models, the parameters were perturbed by additive centered Gaussian noise of variance σ^2 , followed by clippings of h_i s to non-negative values.

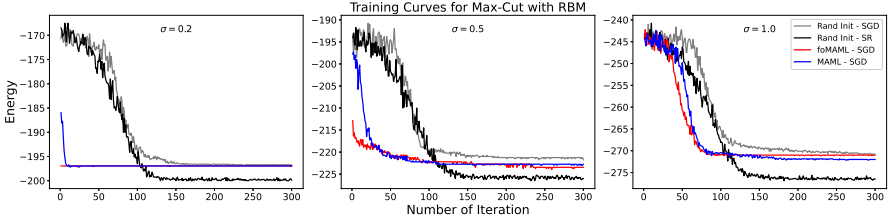


Fig. 1 Results for RBM on $\mathcal{T}_\sigma$ with $\sigma = 0.2, 0.5, 1.0$ from left to right, where the base Hamiltonian is Max-Cut-49. Each curve is the average of the training curves over 8 testing tasks randomly sampled from $\mathcal{T}_\sigma$. The learning rates for outer and inner loops are 0.002 and 0.005, respectively. Initializations from foMAML and MAML discover solutions that outperform random initialization using far fewer iterations with SGD, on problems with diagonal matrix ensembles. However, SR (stochastic reconfiguration) outperforms SGD in the long run. On the other hand, the convergence of foMAML and MAML goes slower as σ increases, indicating that task adaptation becomes more difficult for task distributions that are more complex.

4.2 Architectures and hyper-parameters

Network architectures are chosen to be either restricted Boltzmann machine (RBM) or convolutional neural network (CNN). Carleo *et al.* (2017) [4] proposed RBM for the transverse field Ising model and the anti-ferromagnetic Heisenberg model, taking the one-dimensional state as input and outputs the logarithmic probability amplitude. Besides, RBM accommodates input states with higher dimensional structures by flattening them down to a single dimension. For Ising models with local geometric structures, it’s also natural to use convolutions as local operators to process the inputs. In this paper, we consider CNN models with one layer of convolution followed by a fully connected layer. More architectural details are deferred to the appendices.

Each iteration of the meta-learning loop involves independently sampling a batch of 16 tasks from the task distribution $\mathcal{T}_\sigma$, parametrized by σ . During testing, 8 testing tasks are sampled from $\mathcal{T}_\sigma$ and fixed for evaluation purposes. The inner loop used 1 iteration of vanilla SGD with batch size 1024, while the outer loop training used 50 iterations of vanilla SGD with batch size 16. The learning rates for outer and inner loops depend on the problem type and model architecture.

4.3 Analysis of results

In Figure 1, we train separate RBM models on task distributions $\mathcal{T}_\sigma$ ($\sigma = 0.2, 0.5, 1.0$) with initializations from MAML and foMAML methods, and plot the training curves of the models for 300 iterations. The training curves of randomly initialized models with SGD and stochastic reconfiguration are also plotted for comparison. Our result shows that models initialized using MAML discover solutions that outperform SGD in the long run using far fewer iterations than SGD, consistent with the expectation that MAML initializations perform well after performing only a few iterations of gradient descent. Despite accelerated convergence, MAML was found to converge to suboptimal local

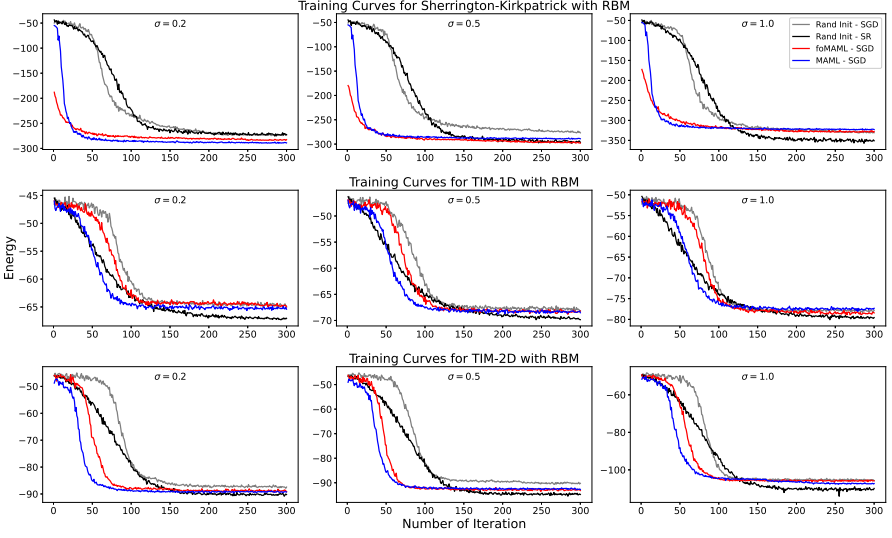


Fig. 2 Results for RBM on $\mathcal{T}_\sigma$ with $\sigma = 0.2, 0.5, 1.0$ from left to right, where the base Hamiltonians are Sherrington-Kirkpatrick model and TIM (transverse field Ising model) with 49 sites. Each curve is the average of the training curves over 8 testing tasks randomly sampled from $\mathcal{T}_\sigma$. The learning rates for outer and inner loops are 0.002 and 0.005, respectively. Initializations from foMAML and MAML discover solutions that outperform random initialization using far fewer iterations with SGD, on problems with sparse non-diagonal matrix ensembles.

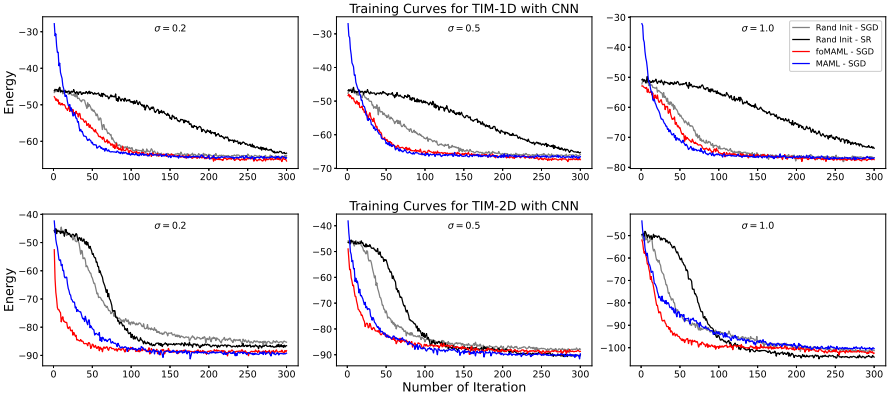


Fig. 3 Results for CNN on $\mathcal{T}_\sigma$ with $\sigma = 0.2, 0.5, 1.0$ from left to right, where the base Hamiltonians are TIM (transverse field Ising model) on 1D lattices with 49 sites and 2D lattices with 7×7 sites. Each curve is the average of the training curves over 8 testing tasks randomly sampled from $\mathcal{T}_\sigma$. The learning rates for outer and inner loops are 0.005 and 0.01, respectively. Initializations from foMAML and MAML discover solutions that outperform random initialization using far fewer iterations with SGD and CNN. On the other hand, the performance of CNN is in general comparable with RBM.

minima in this case compared to the stochastic reconfiguration (initialized from random). This observation reiterates the importance of the Riemannian

geometry to the success of this optimization problem. On the other hand, the convergence of MAML goes slower as σ increases, indicating that task adaptation becomes more difficult for task distributions that are more complex. In Figure 2, we follow similar protocols with the experiments conducted for MaxCut, but switch the base matrix to Sherrington-Kirkpatrick model, and its geometrically local 1D and 2D versions. Similar results are observed for these problems, where MAML can discover solutions that outperform random initialization using far fewer iterations with SGD, and the performance is competitive with that of stochastic reconfiguration training from scratch. This implies that MAML is robust with respect to the choice of Hamiltonian type. In Figure 3, we switch the architecture from RBM to convolutional neural networks that take into account the geometric information of the lattice. The advantage of MAML persists in disordered but geometrically local environments, and MAML performance is largely on par with stochastic reconfiguration. This experiment indicates that MAML is robust with respect to both the choice of the model architecture.

We conclude from our experiments that single-task learners are consistently slow to converge and the convergence remains slow when using the stochastic reconfiguration method. In some cases, for example, the experiments in Figure 3, meta-learners achieved the same or better long-term energy as single-task learners trained with stochastic reconfiguration, at significantly decreased computational cost. In addition, foMAML and MAML exhibit overall better performance on task distribution $\mathcal{T}$ with smaller strength of disorder factor σ , in comparison with the random weight initialization counterparts. This is expected as larger σ corresponds to a sparser task distribution where the samples are less related. The improvements from MAML initializations are robust with respect to the choice of Hamiltonian types and model architectures.

5 Discussion and Future Directions

The experimental results for various matrix ensembles indicate that MAML effectively solves meta-VMC by accelerating training and [improving energy](#). While the Max-Cut problem is exactly solvable for the graph sizes considered here (e.g., by the Branch and Bound method [27]), the experiment illuminates the importance of Riemannian geometry to the success of the algorithm in some learning environments. It would be interesting to investigate if similar findings impact the conclusions of policy-gradient based Meta-RL. [Since stochastic Riemannian optimization appears to be instrumental to the success of VMC in some situations, it would interesting to explore preconditioned task adaptation operators as advocated in \[28\]](#). In the case of quantum spin systems the advantage offered by SR optimization is less significant, while MAML maintains the lead compared to SGD both in terms of acceleration and energy. The reduction in performance with increasing disorder is expected on general grounds, since MAML has less opportunity to exploit regularities in the task distribution. We speculate that the geometrically local models provide further

opportunities to exploit task regularity, explaining the improved performance of MAML relative to SR in these environments. The models investigated all have non-positive off-diagonal entries, which is a simplifying assumption and can be relaxed, paving the way to investigating matrix ensembles relevant to electronic structure. The ideas presented in this paper naturally extend also to variational quantum algorithms (VQAs) such as the variational quantum eigensolver. The key difference in the case of VQAs is that the denominator in the Rayleigh quotient (1) is normalized $\langle \psi_\theta, \psi_\theta \rangle = 1$ and stochastic estimation of the quantum expectation value $\langle \psi_\theta, H_\tau \psi_\theta \rangle$ involves performing measurements in multiple bases if the Hamiltonian contains non-commuting terms. The exploration of meta-VQA and associated learning algorithms is left to future work.

Acknowledgements

Authors gratefully acknowledge support from NSF under grant DMS-2038030.

Disclosure of potential conflicts of interest

The Authors declare that there is no conflict of interest.

References

- [1] Thrun, S., Pratt, L.: Learning to learn. Springer Science and Business Media (2012)
- [2] Finn, C., Abbeel, P., Levine, S.: Model-agnostic meta-learning for fast adaptation of deep networks. In: Proceedings of the 34th International Conference on Machine Learning-Volume 70, pp. 1126–1135 (2017). JMLR. org
- [3] McMillan, W.L.: Ground state of liquid he^4 . Phys. Rev. **138**, 442–451 (1965)
- [4] Carleo, G., Troyer, M.: Solving the quantum many-body problem with artificial neural networks. Science **355**(6325), 602–606 (2017)
- [5] Nomura, Y., Darmawan, A.S., Yamaji, Y., Imada, M.: Restricted boltzmann machine learning for solving strongly correlated quantum systems. Physical Review B **96**(20), 205152 (2017)
- [6] Luo, D., Clark, B.K.: Backflow transformations via neural networks for quantum many-body wave functions. Physical review letters **122**(22), 226401 (2019)

- [7] Stokes, J., Moreno, J.R., Pnevmatikakis, E.A., Carleo, G.: Phases of two-dimensional spinless lattice fermions with first-quantized deep neural-network quantum states. *Physical Review B* **102**(20), 205122 (2020)
- [8] Pfau, D., Spencer, J.S., Matthews, A.G., Foulkes, W.M.C.: Ab initio solution of the many-electron schrödinger equation with deep neural networks. *Physical Review Research* **2**(3), 033429 (2020)
- [9] Gomes, J., McKiernan, K.A., Eastman, P., Pande, V.S.: Classical quantum optimization with neural network quantum states. *arXiv preprint arXiv:1910.10675* (2019)
- [10] Zhao, T., Carleo, G., Stokes, J., Veerapaneni, S.: Natural evolution strategies and variational Monte Carlo. *Machine Learning: Science and Technology* **2**(2), 02–01 (2020)
- [11] Wierstra, D., Schaul, T., Glasmachers, T., Sun, Y., Peters, J., Schmidhuber, J.: Natural evolution strategies. *The Journal of Machine Learning Research* **15**(1), 949–980 (2014)
- [12] Kakade, S.M.: A natural policy gradient. *Advances in neural information processing systems* **14** (2001)
- [13] Sharir, O., Levine, Y., Wies, N., Carleo, G., Shashua, A.: Deep autoregressive models for the efficient variational simulation of many-body quantum systems. *Physical review letters* **124**(2), 020503 (2020)
- [14] Zhao, T., De, S., Chen, B., Stokes, J., Veerapaneni, S.: Overcoming barriers to scalability in variational quantum monte carlo. In: *The International Conference for High Performance Computing, Networking, Storage, and Analysis* (2021)
- [15] Zhao, T., Stokes, J., Knitter, O., Chen, B., Veerapaneni, S.: Meta variational Monte Carlo. *Third Workshop on Machine Learning and the Physical Sciences (NeurIPS 2020)* (2020)
- [16] Williams, R.: Toward a theory of reinforcement-learning connectionist systems. *Technical Report NU-CCS-88-3*, Northeastern University (1988)
- [17] Williams, R.J.: Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning* **8**(3), 229–256 (1992)
- [18] Sorella, S.: Green function monte carlo with stochastic reconfiguration. *Physical Review Letters* **80**(20), 4558–4561 (1998)
- [19] Wilson, M., Stromswold, R., Wudarski, F., Hadfield, S., Tubman, N.M.,

- Rieffel, E.G.: Optimizing quantum heuristics with meta-learning. *Quantum Machine Intelligence* **3**(1), 1–14 (2021)
- [20] Verdon, G., Broughton, M., McClean, J.R., Sung, K.J., Babbush, R., Jiang, Z., Neven, H., Mohseni, M.: Learning to learn with quantum neural networks via classical neural networks. *arXiv preprint arXiv:1907.05415* (2019)
- [21] Andrychowicz, M., Denil, M., Gomez, S., Hoffman, M.W., Pfau, D., Schaul, T., Shillingford, B., De Freitas, N.: Learning to learn by gradient descent by gradient descent. In: *Advances in Neural Information Processing Systems* (2016)
- [22] Caruana, R.: Multitask learning. *Machine learning* **28**(1), 41–75 (1997)
- [23] Baxter, J.: A model of inductive bias learning. *Journal of artificial intelligence research* **12**, 149–198 (2000)
- [24] McCloskey, M., Cohen, N.J.: Catastrophic interference in connectionist networks: The sequential learning problem. In: *Psychology of Learning and Motivation* vol. 24, pp. 109–165 (1989)
- [25] Nichol, A., Achiam, J., Schulman, J.: On first-order meta-learning algorithms. *arXiv preprint arXiv:1803.02999* (2018)
- [26] Fallah, A., Mokhtari, A., Ozdaglar, A.: On the convergence theory of gradient-based model-agnostic meta-learning algorithms. In: *International Conference on Artificial Intelligence and Statistics*, pp. 1082–1092 (2020)
- [27] Rendl, F., Rinaldi, G., Wiegele, A.: Solving max-cut to optimality by intersecting semidefinite and polyhedral relaxations. *Mathematical Programming* **121**(2), 307 (2010)
- [28] Flennerhag, S., Rusu, A.A., Pascanu, R., Visin, F., Yin, H., Hadsell, R.: Meta-learning with warped gradient descent. *arXiv preprint arXiv:1909.00025* (2019)