

Title: Reinforcement Learning assisted Quantum Optimization

Speakers: Matteo Wauters

Series: Machine Learning Initiative

Date: June 05, 2020 - 11:00 AM

URL: <http://pirsa.org/20060019>

Abstract: We propose a reinforcement learning (RL) scheme for feedback quantum control within the quantum approximate optimization algorithm (QAOA). QAOA requires a variational minimization for states constructed by applying a sequence of unitary operators, depending on parameters living in a highly dimensional space. We reformulate such a minimum search as a learning task, where a RL agent chooses the control parameters for the unitaries, given partial information on the system. We show that our RL scheme learns a policy converging to the optimal adiabatic solution for QAOA found by Mbeng et al. arXiv:1906.08948 for the translationally invariant quantum Ising chain. In the presence of disorder, we show that our RL scheme allows the training part to be performed on small samples, and transferred successfully on larger systems. Finally, we discuss QAOA on the p-spin model and how its robustness is enhanced by reinforcement learning. Despite the possibility of finding the ground state with polynomial resources even in the presence of a first order phase transition, local optimizations in the p-spin model suffer from the presence of many minima in the energy landscape. RL helps to find regular solutions that can be generalized to larger systems and make the optimization less sensitive to noise.

References

<https://arxiv.org/abs/2003.07419>

<https://arxiv.org/abs/2004.12323>

# **Reinforcement Learning assisted Quantum Optimization**

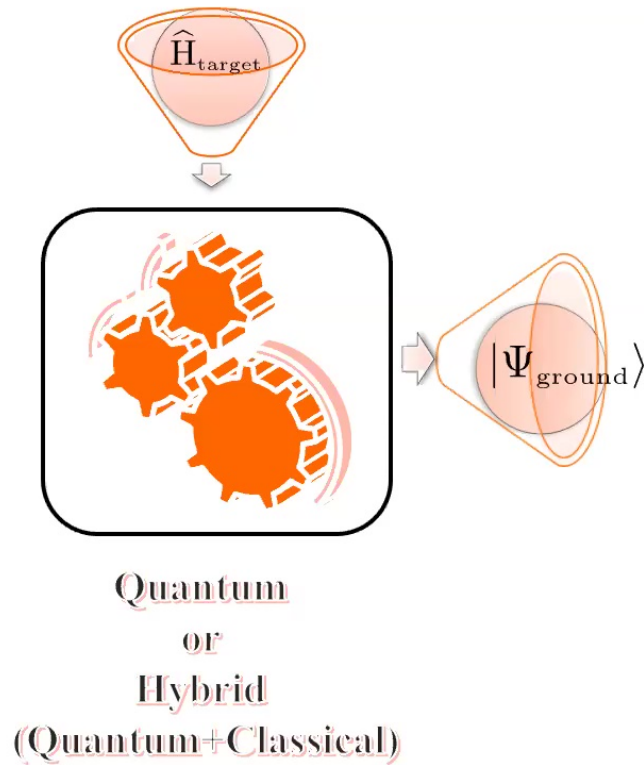
[Preliminary account in: arXiv:2004.12323]



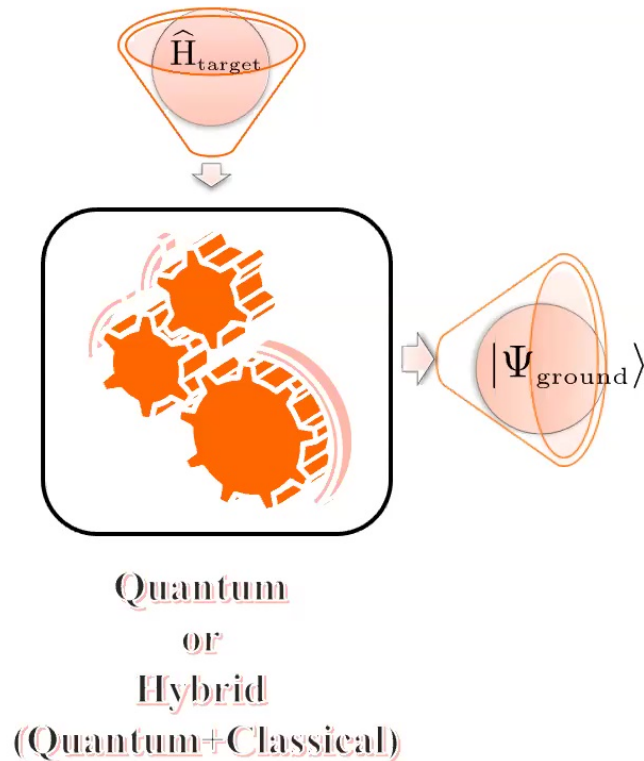
## **SISSA**

**Matteo M. Wauters**  
**SISSA, Trieste**

# The dream of an H-to-Psi machine



# The dream of an H-to-Psi machine



- **Material/Drug design**
- **Condensed Matter problems**
- **Classical Optimization problems (SAT, TSP,...)**

## Algorithms:

- **Adiabatic (QA/AQC)  $\Rightarrow$  continuous evolution**
- **Digital (QAOA, Bang-Bang)  $\Rightarrow$  discrete, stepwise evolution**

# Continuous or discrete?

$$\hat{H}_{\text{target}} = \hat{H}_z = \sum_{\langle i,j \rangle} J_{ij} \hat{\sigma}_i^z \hat{\sigma}_j^z$$

$$\hat{H}_{\text{driving}} = \hat{H}_x = - \sum_i \hat{\sigma}_i^x$$

$$|\psi_0\rangle = |+\rangle = \frac{1}{\sqrt{2^n}} \prod_{i=1}^n (|\uparrow\rangle_i + |\downarrow\rangle_i)$$

**Equal weight quantum superposition**

# Continuous or discrete?

$$\hat{H}_{\text{target}} = \hat{H}_z = \sum_{\langle i,j \rangle} J_{ij} \hat{\sigma}_i^z \hat{\sigma}_j^z$$

$$\hat{H}_{\text{driving}} = \hat{H}_x = - \sum_i \hat{\sigma}_i^x$$

$$|\psi_0\rangle = |+\rangle = \frac{1}{\sqrt{2^n}} \prod_{i=1}^n (|\uparrow\rangle_i + |\downarrow\rangle_i)$$

**Equal weight quantum superposition**

Quantum annealing/adiabatic quantum computing

$$\hat{H}(t) = s(t)\hat{H}_z + (1 - s(t))\hat{H}_x$$

$$s(t) \in [0, 1], \quad s(0) = 0, \quad s(\tau) = 1$$

Total evolution time  $\tau \gg 1$

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = \hat{H}(s(t)) |\psi(t)\rangle$$

Adiabatic theorem  $\Rightarrow |\psi(\tau)\rangle \simeq$  GS of  $\hat{H}_z$

✓ Convergence guaranteed in  $\tau \rightarrow \infty$  limit

# Continuous or discrete?

$$\hat{H}_{\text{target}} = \hat{H}_z = \sum_{\langle i,j \rangle} J_{ij} \hat{\sigma}_i^z \hat{\sigma}_j^z$$

$$\hat{H}_{\text{driving}} = \hat{H}_x = - \sum_i \hat{\sigma}_i^x$$

$$|\psi_0\rangle = |+\rangle = \frac{1}{\sqrt{2^n}} \prod_{i=1}^n (|\uparrow\rangle_i + |\downarrow\rangle_i)$$

**Equal weight quantum superposition**

Quantum annealing/adiabatic quantum computing

$$\hat{H}(t) = s(t)\hat{H}_z + (1 - s(t))\hat{H}_x$$

$$s(t) \in [0, 1], \quad s(0) = 0, \quad s(\tau) = 1$$

Total evolution time  $\tau \gg 1$

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = \hat{H}(s(t)) |\psi(t)\rangle$$

Adiabatic theorem  $\Rightarrow |\psi(\tau)\rangle \simeq \text{GS of } \hat{H}_z$

✓ Convergence guaranteed in  $\tau \rightarrow \infty$  limit

✓ Easy to implement

✗ Sensitive to small gaps

✗ Hard schedule optimization

✗ Time dependent  $|\psi(t)\rangle$  hard to compute

$$|\psi(\tau)\rangle = \text{Texp} \left( -\frac{i}{\hbar} \int_0^\tau dt' \hat{H}(s(t')) \right) |\psi_0\rangle$$

# Continuous or discrete?

Stepwise evolution + Trotter decomposition

$$|\psi(\tau)\rangle = \text{Texp} \left( -\frac{i}{\hbar} \int_0^\tau dt' \hat{H}(s(t')) \right) |\psi_0\rangle \quad \longrightarrow \quad |\psi(\tau)\rangle_{\text{step}} = \prod_{m=1}^{\leftarrow P} e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} |\psi_0\rangle$$
$$e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} \longrightarrow e^{-i\beta_m \hat{H}_x} e^{-i\gamma_m \hat{H}_z} + O(\Delta t_m^2)$$

# Continuous or discrete?

Stepwise evolution + Trotter decomposition

$$|\psi(\tau)\rangle = \text{Texp} \left( -\frac{i}{\hbar} \int_0^\tau dt' \hat{H}(s(t')) \right) |\psi_0\rangle \longrightarrow |\psi(\tau)\rangle_{\text{step}} = \prod_{m=1}^{\leftarrow P} e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} |\psi_0\rangle$$

$$e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} \longrightarrow e^{-i\beta_m \hat{H}_x} e^{-i\gamma_m \hat{H}_z} + O(\Delta t_m^2)$$

digital-QA: from  $s(t)$  to  $s_m$

$$\begin{cases} \gamma_m = s_m \frac{\Delta t_m}{\hbar} \\ \beta_m = (1 - s_m) \frac{\Delta t_m}{\hbar} \end{cases} \longrightarrow \begin{cases} \frac{\Delta t_m}{\hbar} = \gamma_m + \beta_m \\ s_m = \frac{\gamma_m}{\gamma_m + \beta_m} \end{cases}$$

QAOA:  $\gamma_m$  and  $\beta_m$  are variational parameters to optimize with classical methods.

Minimize  $\langle \psi(\gamma, \beta) | \hat{H}_z | \psi(\gamma, \beta) \rangle \longrightarrow \gamma^*, \beta^*$

QAOA

✓ Insensitive to small gaps

# Continuous or discrete?

Stepwise evolution + Trotter decomposition

$$|\psi(\tau)\rangle = \text{Texp} \left( -\frac{i}{\hbar} \int_0^\tau dt' \hat{H}(s(t')) \right) |\psi_0\rangle \longrightarrow |\psi(\tau)\rangle_{\text{step}} = \prod_{m=1}^{\leftarrow P} e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} |\psi_0\rangle$$

$$e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} \longrightarrow e^{-i\beta_m \hat{H}_x} e^{-i\gamma_m \hat{H}_z} + O(\Delta t_m^2)$$

digital-QA: from  $s(t)$  to  $s_m$

$$\begin{cases} \gamma_m = s_m \frac{\Delta t_m}{\hbar} \\ \beta_m = (1 - s_m) \frac{\Delta t_m}{\hbar} \end{cases} \longrightarrow \begin{cases} \frac{\Delta t_m}{\hbar} = \gamma_m + \beta_m \\ s_m = \frac{\gamma_m}{\gamma_m + \beta_m} \end{cases}$$

QAOA:  $\gamma_m$  and  $\beta_m$  are variational parameters to optimize with classical methods.

Minimize  $\langle \psi(\gamma, \beta) | \hat{H}_z | \psi(\gamma, \beta) \rangle \longrightarrow \gamma^*, \beta^*$

QAOA

- ✓ Insensitive to small gaps
- ✓ Easier schedule control (also dQA)
- ✓ More general than dQA
- ✗ Local minima in parameter space
- ✗ Requires many measures for local optimization

# Continuous or discrete?

Stepwise evolution + Trotter decomposition

$$|\psi(\tau)\rangle = \text{Texp} \left( -\frac{i}{\hbar} \int_0^\tau dt' \hat{H}(s(t')) \right) |\psi_0\rangle \longrightarrow |\psi(\tau)\rangle_{\text{step}} = \prod_{m=1}^{\leftarrow P} e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} |\psi_0\rangle$$

$$e^{-\frac{i}{\hbar} \hat{H}(s_m) \Delta t_m} \longrightarrow e^{-i\beta_m \hat{H}_x} e^{-i\gamma_m \hat{H}_z} + O(\Delta t_m^2)$$

digital-QA: from  $s(t)$  to  $s_m$

$$\begin{cases} \gamma_m = s_m \frac{\Delta t_m}{\hbar} \\ \beta_m = (1 - s_m) \frac{\Delta t_m}{\hbar} \end{cases} \longrightarrow \begin{cases} \frac{\Delta t_m}{\hbar} = \gamma_m + \beta_m \\ s_m = \frac{\gamma_m}{\gamma_m + \beta_m} \end{cases}$$

QAOA:  $\gamma_m$  and  $\beta_m$  are variational parameters to optimize with classical methods.

Minimize  $\langle \psi(\gamma, \beta) | \hat{H}_z | \psi(\gamma, \beta) \rangle \longrightarrow \gamma^*, \beta^*$

QAOA

- ✓ Insensitive to small gaps
- ✓ Easier schedule control (also dQA)
- ✓ More general than dQA
- ✗ Local minima in parameter space
- ✗ Requires many measures for local optimization
- ✗ Hard to improve systematically

# Merging dQA and QAOA

- QAOA > dQA
- Exist optimal bang-gang schedule  $s(t)$ 
  - ❖ Belong to QAOA *Ansatz*
  - ❖ Non-smooth, diabatic



PHYSICAL REVIEW X 7, 021027 (2017)

Optimizing Variational Quantum Algorithms Using Pontryagin's Minimum Principle

Zhi-Cheng Yang,<sup>1</sup> Armin Rahmani,<sup>2,3</sup> Alireza Shabani,<sup>4</sup> Hartmut Neven,<sup>4</sup> and Claudio Chamon<sup>1</sup>

# Merging dQA and QAOA

- QAOA > dQA
- Exist optimal bang-gang schedule  $s(t)$ 
  - ❖ Belong to QAOA *Ansatz*
  - ❖ Non-smooth, diabatic
- Iterative optimization of QAOA
  - ❖ Optimal dQA solution
  - ❖ Adiabatic



PHYSICAL REVIEW X 7, 021027 (2017)

Optimizing Variational Quantum Algorithms Using Pontryagin's Minimum Principle

Zhi-Cheng Yang,<sup>1</sup> Armin Rahmani,<sup>2,3</sup> Alireza Shabani,<sup>4</sup> Hartmut Neven,<sup>4</sup> and Claudio Chamón<sup>1</sup>

arXiv:1911.12259

Optimal quantum control with digitized Quantum Annealing

Glen Bigan Mbeng,<sup>1,2</sup> Rosario Fazio,<sup>3,4</sup> and Giuseppe E. Santoro<sup>1,3,5</sup>

<sup>1</sup>SISSA, Via Bonomea 265, I-34136 Trieste, Italy

<sup>2</sup>INFN, Sezione di Trieste, I-34136 Trieste, Italy

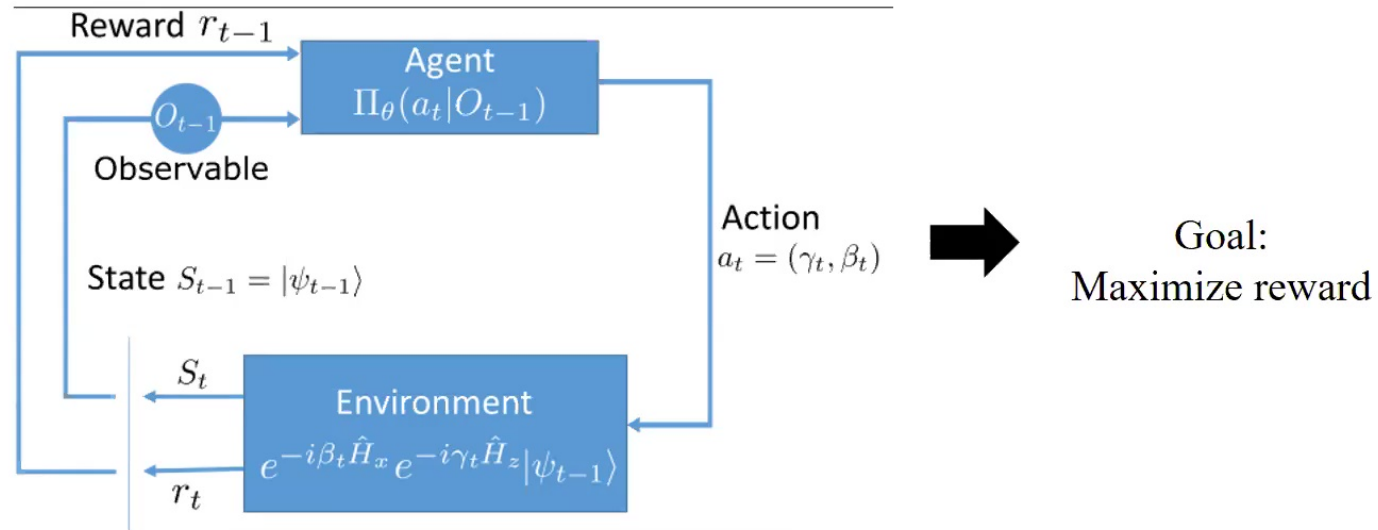
<sup>3</sup>Abdus Salam ICTP, Strada Costiera 11, 34151 Trieste, Italy

<sup>4</sup>Dipartimento di Fisica, Università di Napoli "Federico II", Monte S. Angelo, I-80126 Napoli, Italy

<sup>5</sup>CNR-IOM Democritos National Simulation Center, Via Bonomea 265, I-34136 Trieste, Italy

**Our goal:** Reinforcement Learning can find smooth optimal schedules

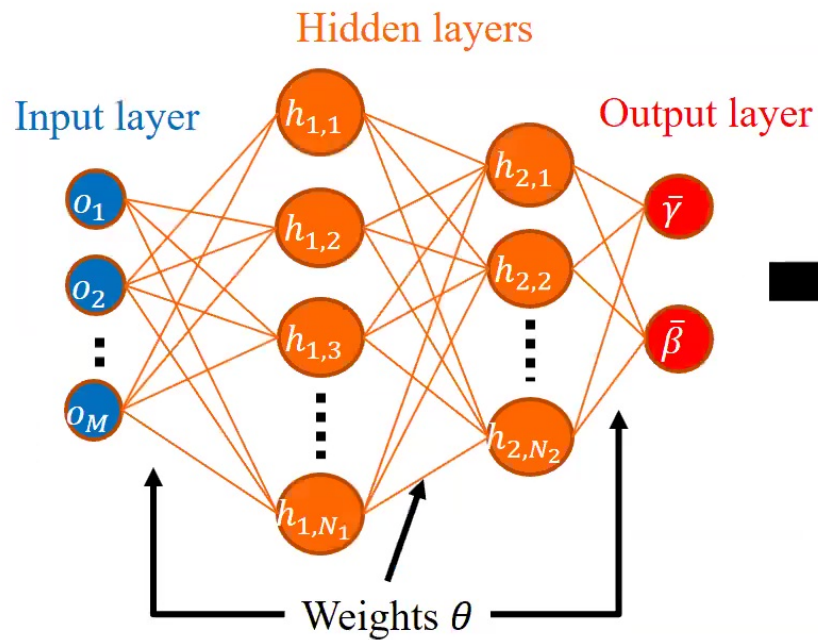
# QAOA as a RL process



At each step  $t = 1, \dots, P$

- Measure  $O_{t-1} = \langle \psi_{t-1} | \hat{O} | \psi_{t-1} \rangle$
- Compute action  $a_t = (\gamma_t, \beta_t)$  from policy  $\Pi_{\theta}(a | O)$
- Update state  $|\psi_t\rangle = e^{-i\beta_t \hat{H}_x} e^{-i\gamma_t \hat{H}_z} |\psi_{t-1}\rangle$
- Compute new reward  $r_t$ : measure quality of variational Ansatz

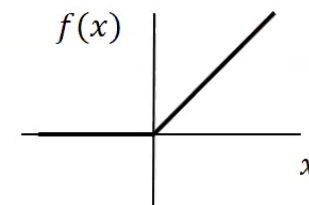
# Choosing the policy



Gaussian policy

$$\Pi_{\theta}(\gamma|o) = e^{-\frac{(\gamma - \bar{\gamma})^2}{2\sigma_{\gamma}^2}}$$

$$\Pi_{\theta}(\beta|o) = e^{-\frac{(\beta - \bar{\beta})^2}{2\sigma_{\beta}^2}}$$

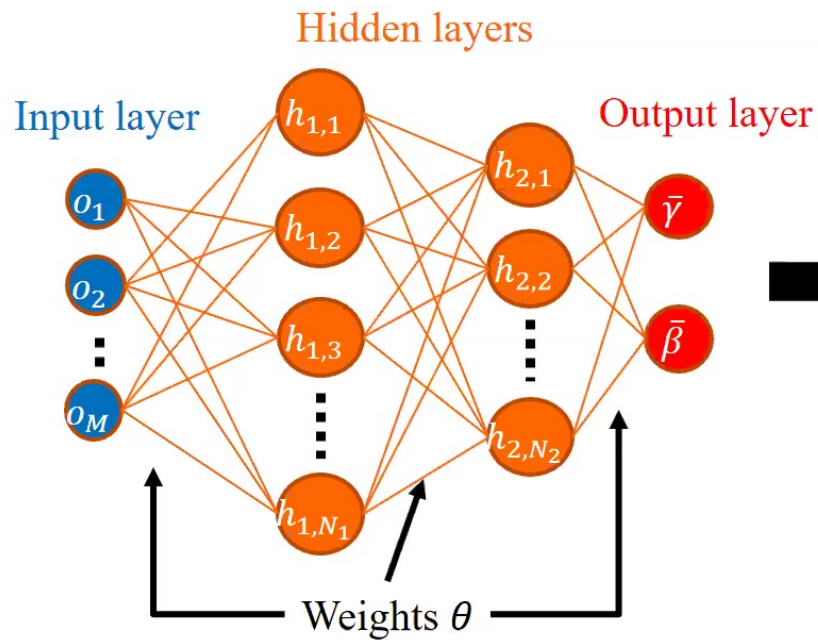


Our choice:

$$O_t = (\langle \psi_t | \hat{H}_x | \psi_t \rangle, \langle \psi_t | \hat{H}_z | \psi_t \rangle)$$

NN: 32,16 fully connected hidden, ReLu activation function

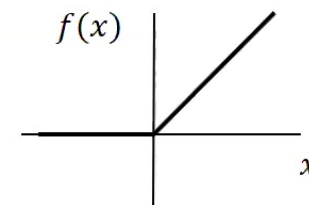
# Choosing the policy



Gaussian policy

$$\Pi_{\theta}(\gamma|o) = e^{-\frac{(\gamma - \bar{\gamma})^2}{2\sigma_{\gamma}^2}}$$

$$\Pi_{\theta}(\beta|o) = e^{-\frac{(\beta - \bar{\beta})^2}{2\sigma_{\beta}^2}}$$

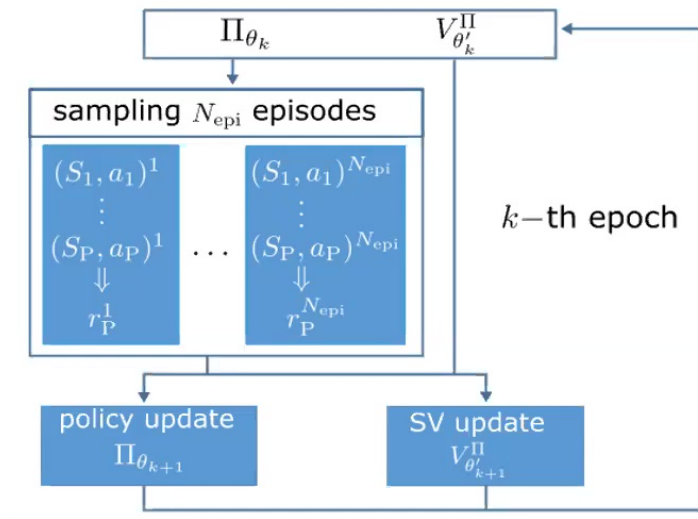


Our choice:

$$O_t = (\langle \psi_t | \hat{H}_x | \psi_t \rangle, \langle \psi_t | \hat{H}_z | \psi_t \rangle)$$

NN: 32,16 fully connected hidden, ReLu activation function

# Learning process



Proximal Policy Optimization (PPO)  
algorithm from OpenAI Spinningup  
library

$k = 1, \dots, N_{\text{epo}}$  epochs

In each epoch:

- sample  $N_{\text{epi}}$  digital evolutions (episodes)
- Update policy evaluation  
 $V_{\theta'}^{\Pi}(o) = \mathbb{E}^{\Pi}(r_P | o_t = o)$
- Update policy  $\Pi_{\theta}$

## Learning parameters

- $N_{\text{epi}} = 100$  episodes per epoch
- $N_{\text{epo}} = 1024$  epochs
- Both policy  $\Pi_{\theta}(a|o)$  and state-value function  $V_{\theta'}^{\Pi}(o)$  parametrized by 32,16 NN
- Reward  $r_t = -\delta_{t,P} \langle \psi_t | \hat{H}_T | \psi_t \rangle$

# Technical interlude: PPO details

NN updates:

- Compute gradients  $\partial_{\theta} \Pi_{\theta}(a|o), \partial_{\theta'} V_{\theta'}^{\Pi}(o) \Rightarrow$  backpropagation (chain rule)
- Decide what to do with gradient  $\Rightarrow$  PPO algorithm

Value function  $V^{\Pi}(o) = \mathbb{E}(R|o_0 = o)$

Expected future reward with policy  $\Pi$ , starting from  $o$

Update: minimize  $|V^{\Pi}(o) - V_{\theta'}^{\Pi}(o)|$

$R = r_1 + \lambda r_2 + \lambda^2 r_3 + \dots$   
 $\lambda \leq 1$  discount factor  
 $V^{\Pi}(o)$  **real** VF, estimated from experience  
 $V_{\theta'}^{\Pi}(o)$  **approximated** by NN

Policy  $\Pi_{\theta}(a|o)$  : probability of taking action  $a$  when observing  $o$

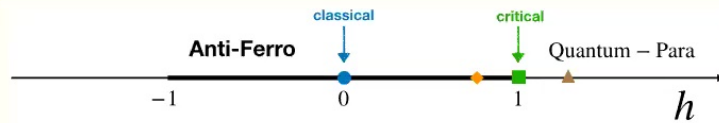
Given action  $a : o \rightarrow o', r \Rightarrow$  define **advantage**  $A^{\Pi}(o, a) = (r + V^{\Pi}(o')) - V^{\Pi}(o)$

Update: maximize  $\frac{\Pi_{\theta}(a|o)}{\Pi_{\theta_k}(a|o)} A^{\Pi_{\theta_k}}(o, a)$

$\theta_k$  are the old parameters  
Update is **clipped** to control maximum change

# Quantum Ising chain

[... on an N-ring with PBC]



$$\hat{H}_T = \hat{H}_z + h\hat{H}_x$$

$$\hat{H}_z = \sum_{j=1}^N \hat{\sigma}_j^z \hat{\sigma}_{j+1}^z$$

$$\hat{H}_x = - \sum_{j=1}^N \hat{\sigma}_j^x$$

$$|\psi_P(\gamma, \beta)\rangle = \hat{U}_P \cdots \hat{U}_1 |\psi_0\rangle$$

$$\hat{U}_m = e^{-i\beta_m \hat{H}_x} e^{-i\gamma_m \hat{H}_z}$$

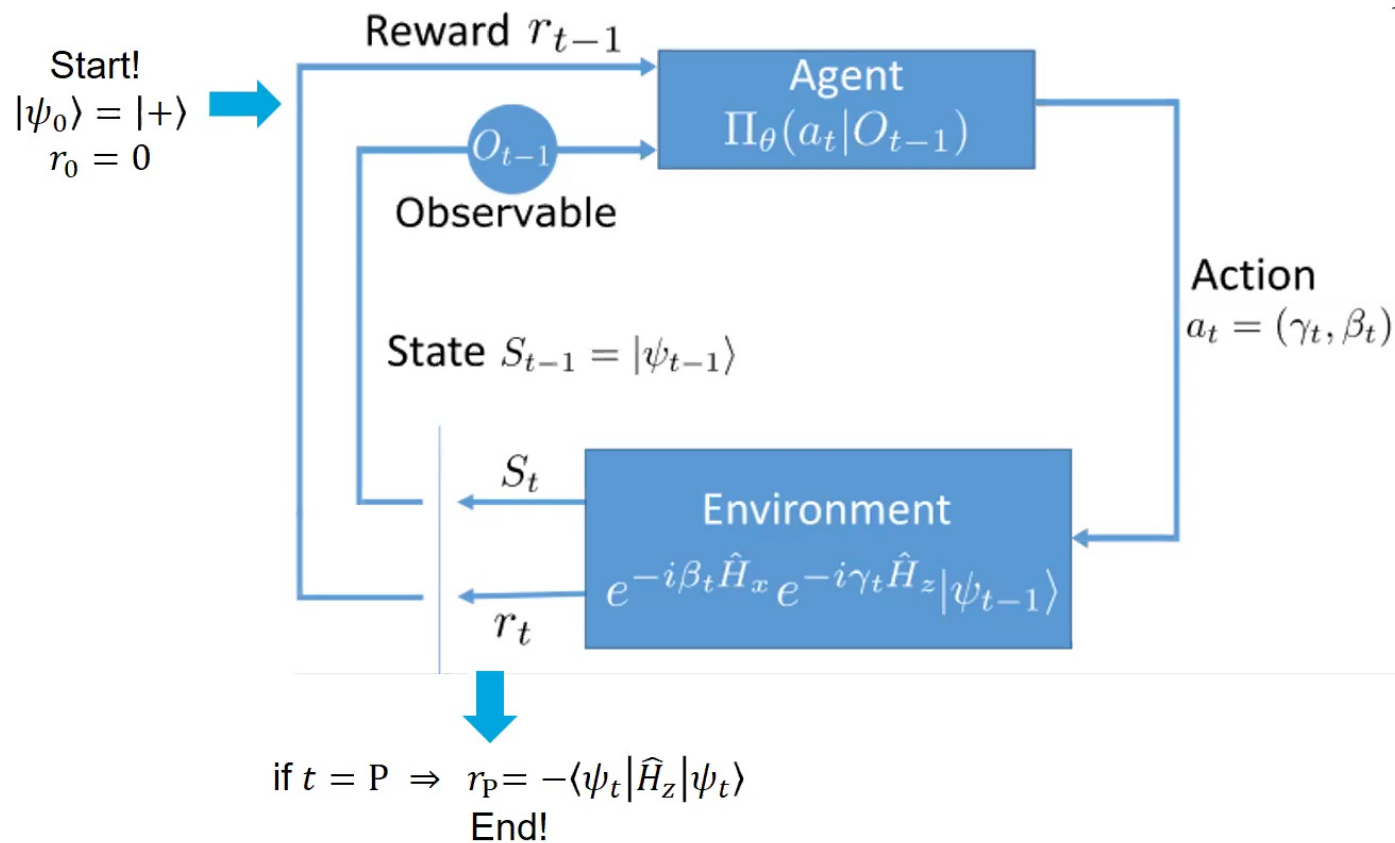
$$E_P(\gamma, \beta) = \langle \psi_P(\gamma, \beta) | \hat{H}_T | \psi_P(\gamma, \beta) \rangle$$

$$\text{Reward } r_P = -E_P(\gamma, \beta)$$

## Residual energy

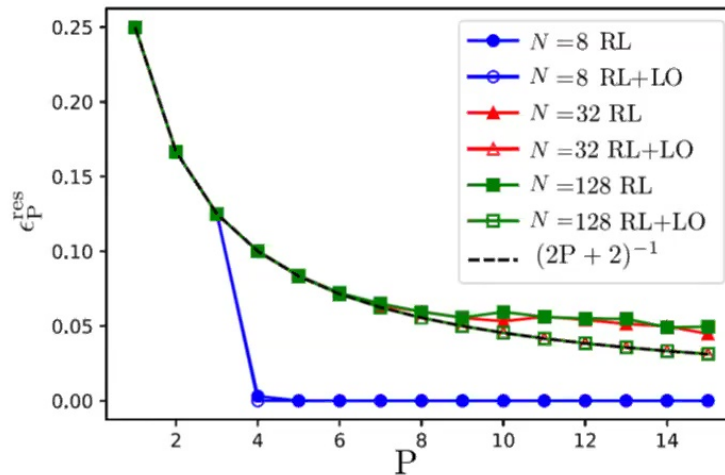
$$\epsilon_P^{\text{res}}(\gamma, \beta) = \frac{E_P(\gamma, \beta) - E_{\text{gs}}}{E_{\text{max}} - E_{\text{gs}}} \in [0, 1]$$

# Recap: QAOA as a RL process



# Quantum Ising chain

[... on an N-ring with PBC]



Test of the trained agent

- $P \leq 7$ : RL alone replicates QAOA solution.
- $P > 7$ : RL follows QAOA solution but needs local optimization for exact match

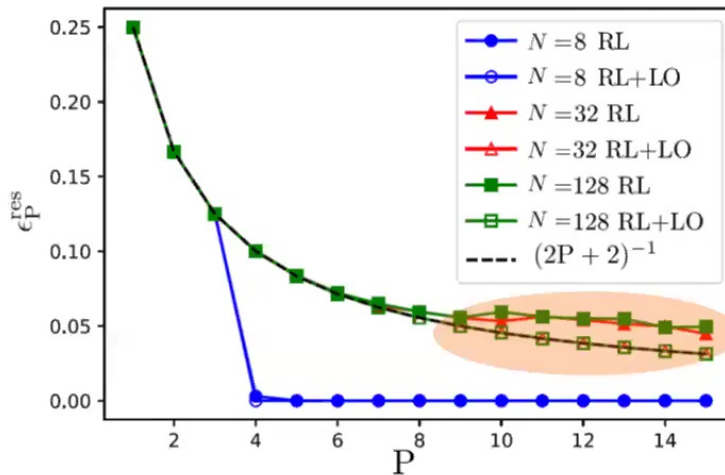
Lower bound for TFIM in arbitrary dimension, exact in 1d

$$\epsilon_P^{\text{res}}(\gamma, \beta) \geq \begin{cases} \frac{1}{2P+2} & 2P < N \\ 0 & 2P \geq N \end{cases}$$

Mbeng et al. arXiv:1906.08948

# Quantum Ising chain

[... on an N-ring with PBC]



Test of the trained agent

- $P \leq 7$ : RL alone replicates QAOA solution.
- $P > 7$ : RL follows QAOA solution but needs local optimization for exact match

Lower bound for TFIM in arbitrary dimension, exact in 1d

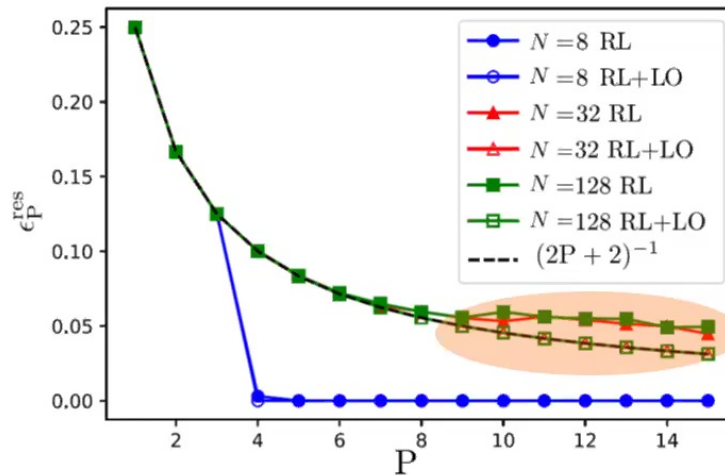
$$\epsilon_P^{\text{res}}(\gamma, \beta) \geq \begin{cases} \frac{1}{2P+2} & 2P < N \\ 0 & 2P \geq N \end{cases}$$

Mbeng et al. arXiv:1906.08948

- Convergence in training epochs not reached
- Need of fine tuning of training parameters
- Stochastic policy  $\Rightarrow$  larger trajectory dispersion for longer episodes

# Quantum Ising chain

[... on an N-ring with PBC]



Test of the trained agent

- $P \leq 7$ : RL alone replicates QAOA solution.
- $P > 7$ : RL follows QAOA solution but needs local optimization for exact match

Lower bound for TFIM in arbitrary dimension, exact in 1d

$$\epsilon_P^{\text{res}}(\gamma, \beta) \geq \begin{cases} \frac{1}{2P+2} & 2P < N \\ 0 & 2P \geq N \end{cases}$$

Mbeng et al. arXiv:1906.08948

- Convergence in training epochs not reached
- Need of fine tuning of training parameters
- Stochastic policy  $\Rightarrow$  larger trajectory dispersion for longer episodes

# Recap: QAOA and dQA

➤ QA/AQC:  $\hat{H}(t) = s(t)\hat{H}_z + (1 - s(t))\hat{H}_x$

❖  $s(t) \in [0,1]$

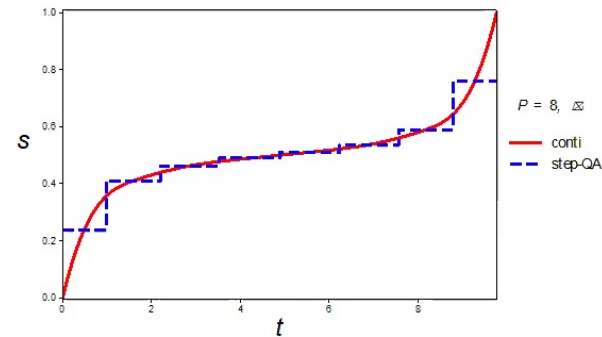
❖ Evolution time  $\tau$

➤ dQA:  $s(t) \rightarrow s_t$

❖  $t = 1, \dots, P$

❖  $e^{-\frac{i}{\hbar}\hat{H}_t\Delta t} \rightarrow e^{-\frac{i}{\hbar}\beta_t\hat{H}_x}e^{-\frac{i}{\hbar}\gamma_t\hat{H}_z}$

❖  $s_t = \frac{\gamma_t}{\gamma_t + \beta_t}$



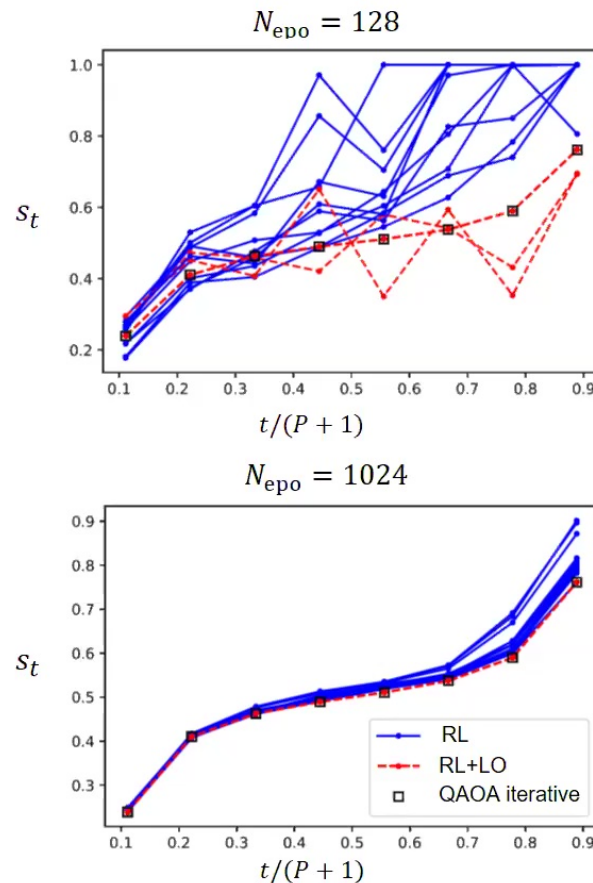
➤ QAOA:  $(\gamma^*, \beta^*) = \arg \min_{(\gamma, \beta)} E_P(\gamma, \beta)$

❖ Classical minimization in 2P-dimensional parameter space

❖  $\exists (\gamma^*, \beta^*)$  such that  $s_t^*$  is the discretization of continuous schedule + adiabatic

# Quantum Ising chain

[... on an N-ring with PBC]



Action choice analysis

$$s_t = \frac{\gamma_t}{\gamma_t + \beta_t}$$

Allows comparison with digital-QA schedule

$$\hat{H}_t = s_t \hat{H}_T + (1 - s_t) \hat{H}_x$$

Training epochs

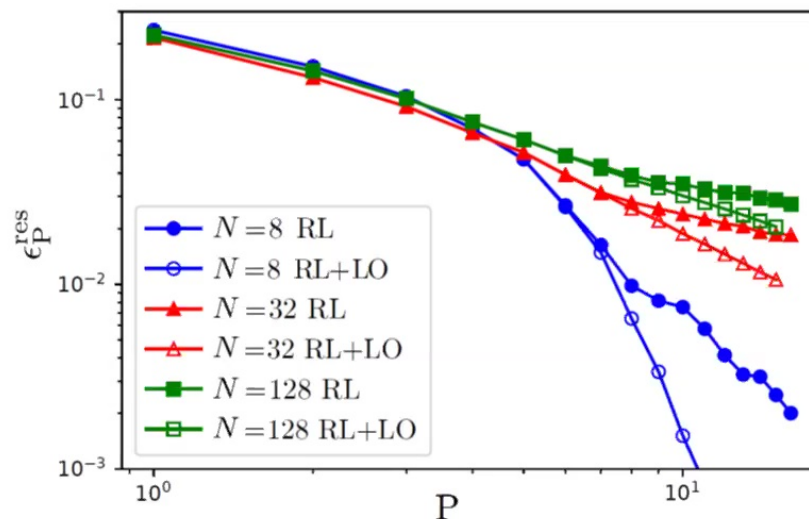
- i. Random action choices
- ii. Large trajectory dispersion. RL+LO fall on the same smooth minimum (QAOA iter.)
- iii. RL converges to QAOA iterative solution

# Quantum Ising chain

With random couplings

$$\hat{H}_z = \sum_{i=1}^N J_i \hat{\sigma}_i^z \hat{\sigma}_{i+1}^z \quad \text{with } J_i \in [0,1], \text{ uniformly distributed}$$

Harder problem for optimization than uniform Ising model



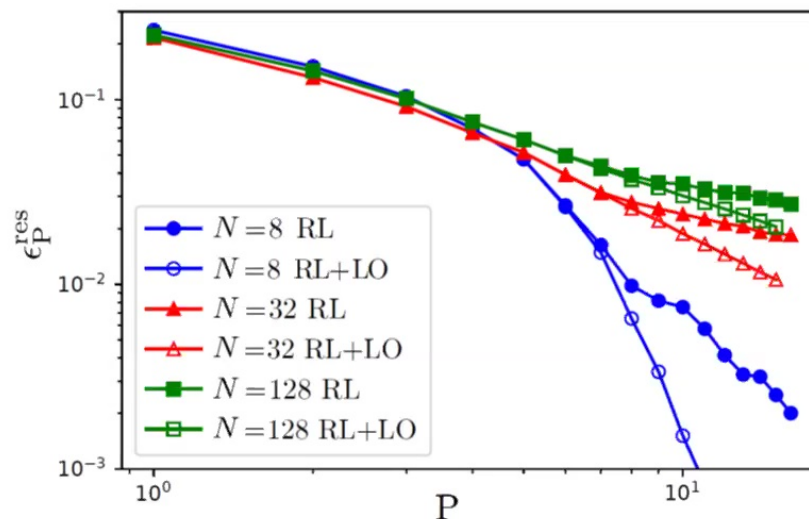
Same phenomenology of uniform TFIM: local optimization leads to lower energy for  $P \geq 8$

# Quantum Ising chain

With random couplings

$$\hat{H}_z = \sum_{i=1}^N J_i \hat{\sigma}_i^z \hat{\sigma}_{i+1}^z \quad \text{with } J_i \in [0,1], \text{ uniformly distributed}$$

Harder problem for optimization than uniform Ising model



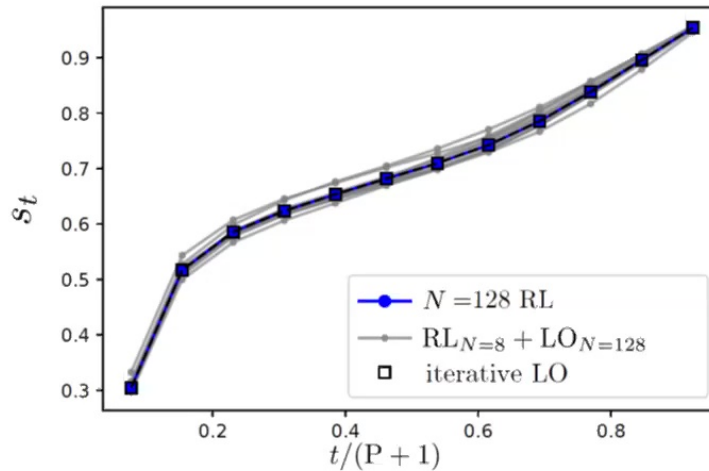
Same phenomenology of uniform TFIM: local optimization leads to lower energy for  $P \geq 8$

Open issues:

- ? No analytical result
- ? Is solution optimal?
- ? Better than QA?

# Quantum Ising chain

With random couplings



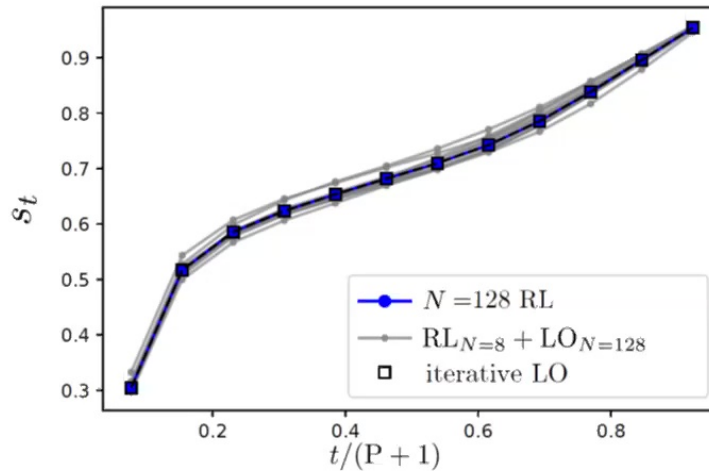
Policy transferability:

- Train a small system ( $N = 8$ ) on single disorder instance
- Get approximate solution  $(\gamma^*, \beta^*)_{N=8}$

Same smooth parameter choice in RL and iterative local optimization

# Quantum Ising chain

With random couplings



Same smooth parameter choice in  
RL and iterative local optimization

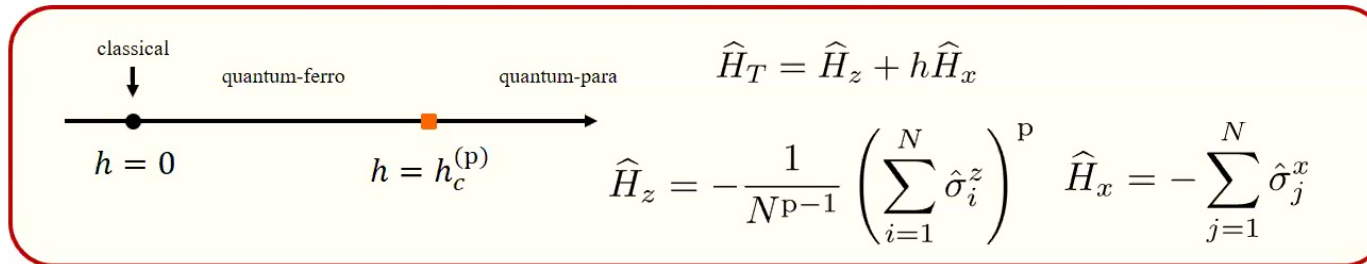
Policy transferability:

- Train a small system ( $N = 8$ ) on single disorder instance
- Get approximate solution  $(\gamma^*, \beta^*)_{N=8}$
- Use  $(\gamma^*, \beta^*)_{N=8}$  as ansatz to initialize a local optimization on larger system ( $N = 128$ ) with any disorder instance



Possible solution to measurement  
problem

# Fully connected Ising model



$p = 2$

- Second order phase transition

$$h_c^{(2)} = 2$$

- Minimum energy gap  $\Delta \sim N^{-\frac{1}{3}}$
- Annealing time  $\tau \sim N^{\frac{2}{3}}$

$p > 2$

- First order phase transition

$$1 < h_c^{(p)} < 2$$

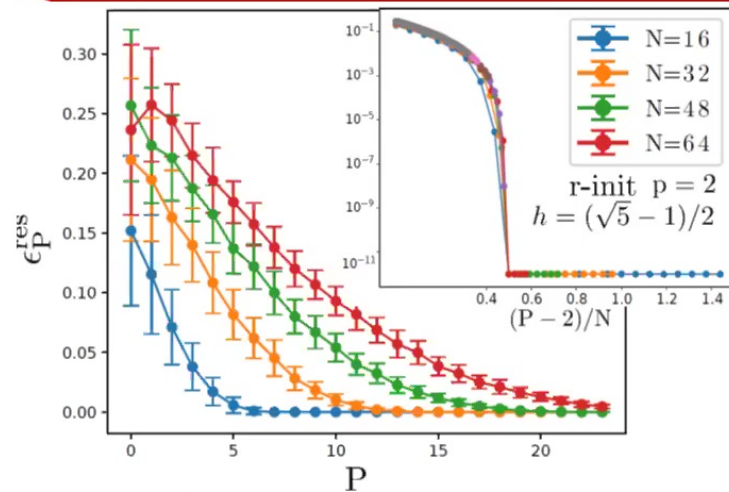
- Minimum energy gap  $\Delta \sim e^{-\alpha_p N}$
- Annealing time  $\tau \sim e^{2\alpha_p N}$



# Fully connected Ising model

arXiv:2003.07419

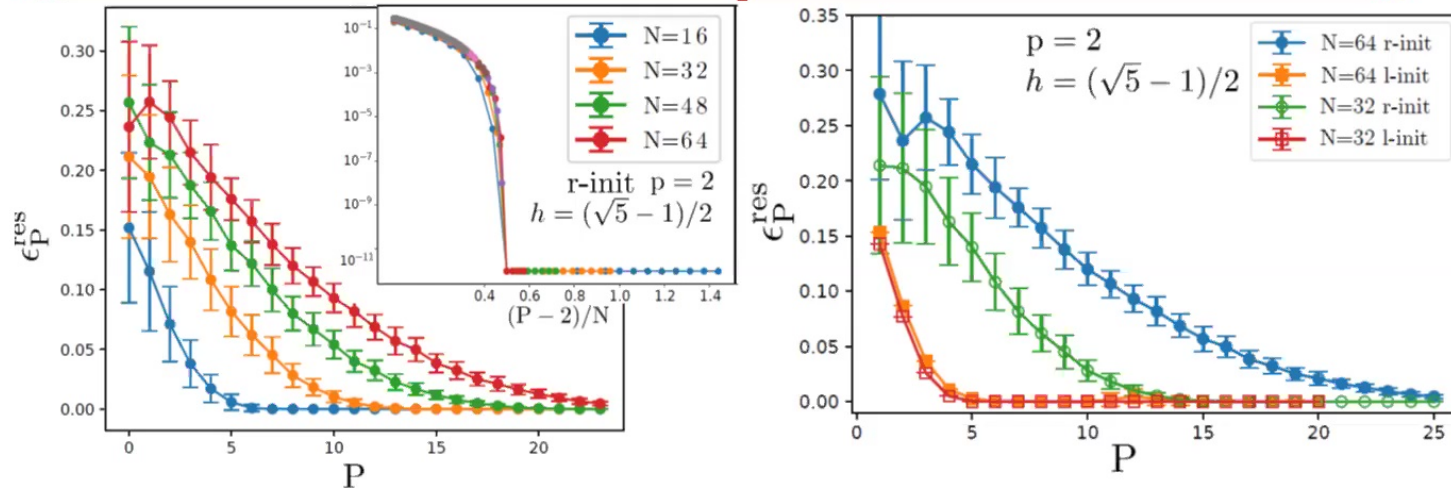
- Results with QAOA (local optimization) only:
- The  $h = 0$  ground state  $|\uparrow\rangle + |\downarrow\rangle$  prepared exactly with  $P=1$  (if  $N$  odd) or  $P=2$  (if  $N$  even) with  $\gamma \sim N^{p-1}$
- Any state can be prepared exactly with  $P = \frac{N}{2} + 2$  ( $p$  even) or  $P = N + 1$  ( $p$  odd).



# Fully connected Ising model

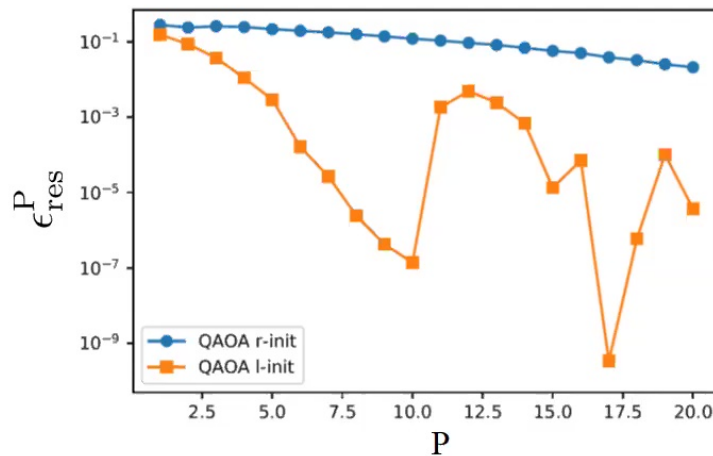
arXiv:2003.07419

- Results with QAOA (local optimization) only:
- The  $h = 0$  ground state  $|\uparrow\rangle + |\downarrow\rangle$  prepared exactly with  $P=1$  (if  $N$  odd) or  $P=2$  (if  $N$  even) with  $\gamma \sim N^{p-1}$
- Any state can be prepared exactly with  $P = \frac{N}{2} + 2$  ( $p$  even) or  $P = N + 1$  ( $p$  odd).
- Exact ground state with polynomial resources.
- Many local minima  $\Rightarrow$  local optimizations depend strongly on initial parameters.
- Hard to find smooth control parameters  $s_t$ .

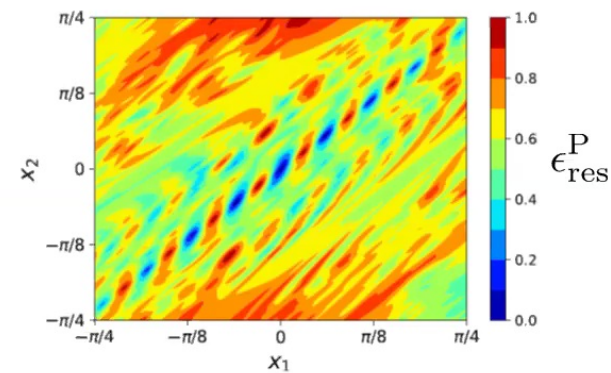


# Fully connected Ising model

$$N = 64, \quad h = \frac{\sqrt{5} - 1}{2}$$



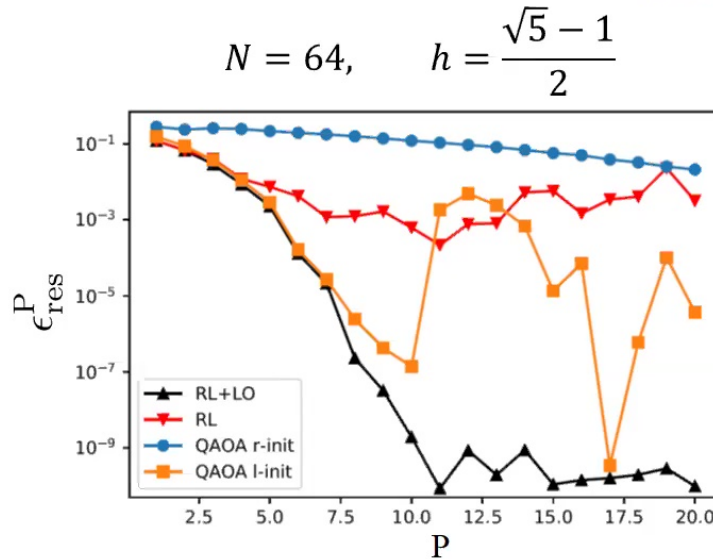
Rugged energy landscape  
Strong dependence from initial  
parameter set



r-init  $(\gamma^0, \beta^0) \in [0, \pi)$

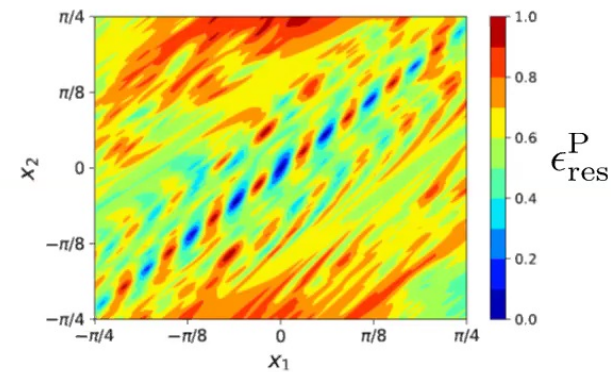
$$\text{l-init} \quad \begin{cases} \gamma_t^0 = \frac{s_t \Delta t_t}{\hbar} \\ \beta_t^0 = \frac{\Delta t_t}{\hbar} (1 - s_t(1 - h)) \end{cases}$$

# Fully connected Ising model



- RL good only for small  $P$
- RL + LO improves and stabilizes results over local optimization alone

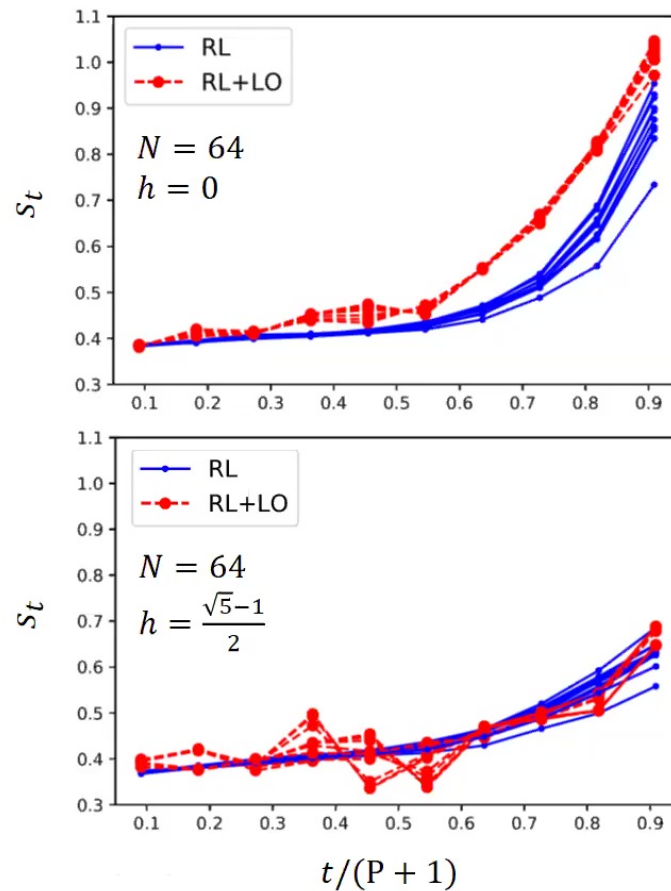
Rugged energy landscape  
Strong dependence from initial parameter set



r-init  $(\gamma^0, \beta^0) \in [0, \pi)$

l-init 
$$\begin{cases} \gamma_t^0 = \frac{s_t \Delta t_t}{\hbar} \\ \beta_t^0 = \frac{\Delta t_t}{\hbar} (1 - s_t(1 - h)) \end{cases}$$

# Fully connected Ising model



Reinforcement learning

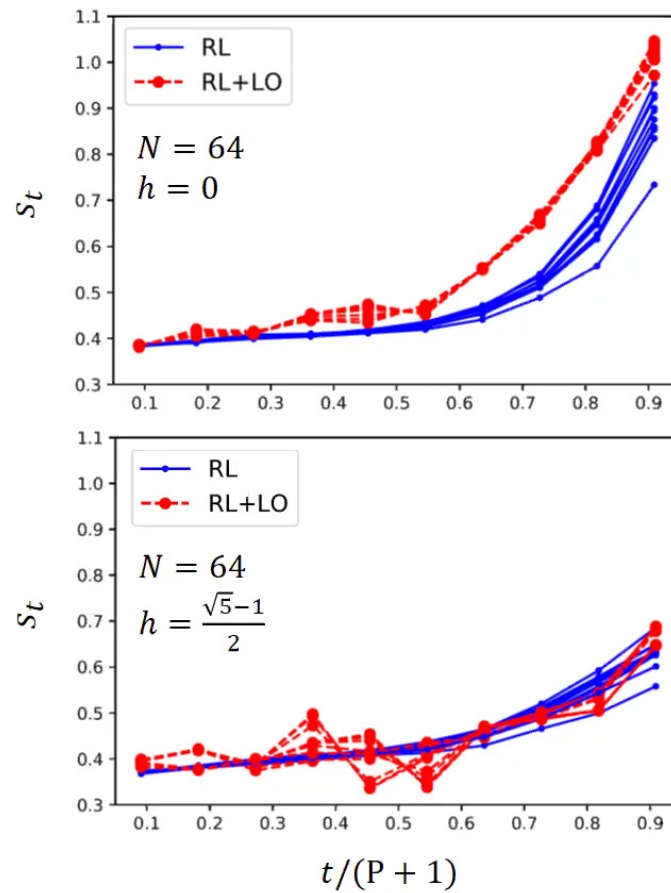
- ✓ suggests smooth action choices
- ✗ Low performance for large P

RL+LO

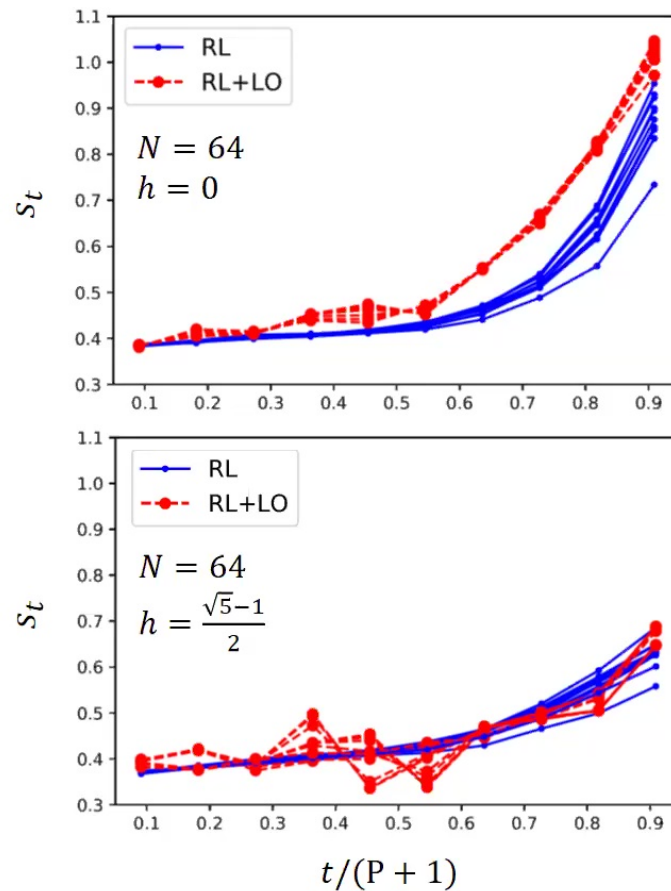
- ✓ dramatically enhances performance
- ✗ Optimized parameters less smooth
- ✗ Different local minima

Is schedule adiabatic?  
Bad local minima due to non-local interactions?

# Fully connected Ising model



# Fully connected Ising model



Reinforcement learning

- ✓ suggests smooth action choices
- ✗ Low performance for large P

RL+LO

- ✓ dramatically enhances performance
- ✗ Optimized parameters less smooth
- ✗ Different local minima

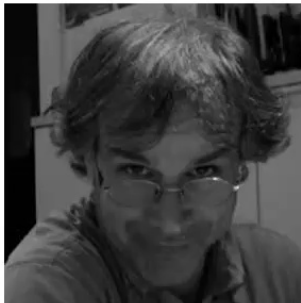
Is schedule adiabatic?  
Bad local minima due to non-local interactions?

# Summary & outlook

- Parameters for QAOA ansatz can be learned with RL approach
  - ✓ Smooth action choices
  - ✓ Few observables needed
  - ✓ Converges to adiabatic optimal schedules
  - ✓ Strategies can be transferred among different system sizes
  - ✗ Requires observation on the system during the evolution
  
- Open issues
  - ? Solve measurement problem: weak measures, ancillary qubits,...?
  - ? Robustness to noise
  - ? Characterize collapse of learned policy during the training
  - ? More complex Hamiltonians (DMRG/MPS?)

# Reinforcement Learning assisted Quantum Optimization

[Preliminary account in: arXiv:2004.12323]



Giuseppe E. Santoro  
The boss



Emanuele Panizon  
The RL guy



Glen B. Mbeng  
The QAOA expert

Thank you!