
What is my quantum computer good for? Quantum capability learning with physics-aware neural networks

Daniel Hothem

Quantum Performance Laboratory
Sandia National Laboratories
Livermore, CA 94550
dhothem@sandia.gov

Ashe Miller

Quantum Performance Laboratory
Sandia National Laboratories
Albuquerque, NM 87185
anmille@sandia.gov

Timothy Proctor

Quantum Performance Laboratory
Sandia National Laboratories
Livermore, CA 94550
tjproct@sandia.gov

Abstract

Quantum computers have the potential to revolutionize diverse fields, including quantum chemistry, materials science, and machine learning. However, contemporary quantum computers experience errors that often cause quantum programs run on them to fail. Until quantum computers can reliably execute large quantum programs, stakeholders will need fast and reliable methods for assessing a quantum computer’s capability—i.e., the programs it can run and how well it can run them. Previously, off-the-shelf neural network architectures have been used to model quantum computers’ capabilities, but with limited success, because these networks fail to learn the complex quantum physics that determines real quantum computers’ errors. We address this shortcoming with a new quantum-physics-aware neural network architecture for learning capability models. Our scalable architecture combines aspects of graph neural networks with efficient approximations to the physics of errors in quantum programs. This approach achieves up to $\sim 50\%$ reductions in mean absolute error on both experimental and simulated data, over state-of-the-art models based on convolutional neural networks, and scales to devices with 100+ qubits.

1 Introduction

Quantum computers have the potential to efficiently solve classically intractable problems in quantum chemistry [Cao et al., 2019], materials science [Rubin et al., 2024], machine learning [Harrow et al., 2009], and cryptography [Shor, 1997]. While contemporary quantum computers are approaching the size and noise levels needed to solve interesting problems [Arute et al., 2019], they are far from being capable of reliably running most useful quantum programs [Proctor et al., 2021a]. Until we build quantum computers capable of executing *any and all* useful and interesting quantum programs, stakeholders will require fast, reliable, and scalable methods for predicting the programs that a given quantum computer can reliably execute.

The task of learning which quantum programs a particular quantum computer can reliably execute is known as *quantum capability learning* [Proctor et al., 2021a]. Quantum capability learning is very

difficult because the number of possible (Markovian) errors plaguing a quantum computer grows exponentially in its size [Blume-Kohout et al., 2022], i.e., in the number of qubits (n) it contains, and errors in a quantum program can combine in difficult-to-predict ways [Proctor et al., 2021a]. Most existing approaches to capability learning restrict themselves to learning how well a quantum computer executes a small set of quantum programs, by running all of those programs and estimating a success metric for each one [Lubinski et al., 2023, Proctor et al., 2024]. While these methods provide insight into a quantum computer’s capability, they are not predictive.

Recently, several groups have proposed building predictive models of a quantum computer’s capability using convolutional neural networks (CNNs) [Elsayed Amer et al., 2022, Hothem et al., 2024b, Vadali et al., 2024, Hothem et al., 2023b] and graph neural networks (GNNs) [Wang et al., 2022]. However, these neural-network-based capability models achieve only modest prediction accuracy when applied to real quantum computers, because they fail to learn the complex physics that determines real quantum computers’ failures [Hothem et al., 2024b].

In this work, we introduce a novel quantum-physics-aware neural network (qpa-NN) architecture for quantum capability learning (Fig. 1). Our approach uses neural networks with GNN-inspired structures to predict the rates of the most physically relevant errors in quantum programs. These predicted error rates are then combined using an efficient approximation to the exact (but exponentially costly) quantum physics formula for how those errors combine to impact a program’s success rate. Our approach leverages the graph structures that encode the physics of how errors’ rates typically depend on both the quantum program being run and how a quantum computer’s qubits are arranged, and it offloads the difficult-to-learn, yet classically tractable task of approximately combining these error rates to predict a circuit’s performance to an already-known function. This enables our qpa-NNs to vastly outperform the state-of-the-art CNNs of Hothem et al. [2024b] on both experimental and simulated data without sacrificing the ability to model large devices of 100+ qubits.

Our qpa-NNs are enabled, in part, by focusing on learning a quantum computer’s capability on high-fidelity quantum programs, which are those programs that a quantum computer correctly executes with high probability. High-fidelity programs are arguably the most interesting programs to study as we care far more about whether a quantum computer successfully executes a program 99% or 90% of the time rather than 1% or 10% of the time.

In a head-to-head comparison, our qpa-NNs achieve a $\sim 50\%$ reduction in mean absolute error (MAE) over the CNNs of Hothem et al. [2024b], on average and on the same experimental datasets. Our qpa-NNs achieve an average $\sim 36\%$ improvement over those CNNs even after fine-tuning those CNNs on the same subset of the training data (high-fidelity programs) used to train our qpa-NNs.

Our qpa-NNs’ improved performance is likely largely due to their improved ability to model the impact of coherent errors on a program’s success rate. Off-the-shelf networks struggle with coherent errors [Hothem et al., 2024b], but our qpa-NNs are designed to model how these errors add up and cancel out, making the qpa-NNs much better predictors in the presence of coherent errors. To verify this, we demonstrate that our qpa-NNs can accurately predict the performance of random circuits run on a hypothetical 4-qubit quantum computer experiencing only coherent errors. Our qpa-NN obtained a $\sim 50\%$ lower MAE than a CNN, averaged across five datasets, and the trained qpa-NN even exhibits moderate performance when making predictions for a different class of circuits (random mirror circuits [Proctor et al., 2021a]) simulated on the same hypothetical 4-qubit quantum computer, i.e., our qpa-NNs display moderate prediction accuracy on out-of-distribution data.

We make the following contributions in our work:

1. We introduce qpa-NNs, a bespoke neural network architecture for modeling the capability of a quantum computer, which outperform state-of-the-art CNN models by $\sim 50\%$ on experimental and simulated data.
2. We use our qpa-NNs to model the capability of a simulated 100-qubit device; the largest ever neural network capability learning demonstration by a factor of two.
3. We demonstrate, for the first time, how to train NNs to predict the process fidelity [Nielsen, 2002] of a circuit, which is the most widely used quantum channel error metric.
4. We provide evidence that the improved performance of our qpa-NNs is partly due to their ability to better model the effect of coherent errors, which are known to be challenging for other state-of-the-art methods.

2 Background

In this section, we review the background in quantum computing necessary to understand this paper. See Nielsen and Chuang [2010] for an in-depth introduction to quantum computing and Blume-Kohout et al. [2022] or Hashim et al. [2024] for a thorough description of the errors in quantum computers.

2.1 Quantum computing

A quantum computer performs computations using qubits, which are two-level systems whose pure states are unit vectors in a complex two-dimensional Hilbert space, $\mathcal{H}$. The pure states of n qubits are unit vectors in $\mathcal{H}^{\otimes n}$. The two orthonormal vectors $|0\rangle$ and $|1\rangle$ that are eigenvectors of the Z Pauli operator are identified as the *computational basis* of $\mathcal{H}$. Errors and noise in real quantum computers mean that they are typically in states ρ that are probabilistic mixtures of pure states.

A quantum computation is performed by running a quantum program, typically known as a *quantum circuit* (see illustration in Fig. 1a). An n -qubit quantum circuit (c) of depth d is defined by a sequence of d layers of logical instructions $\{L_i\}$. Executing c consists of preparing each qubit in $|0\rangle$, applying each L_i , and then measuring each qubit to obtain an n -bit string b . Each layer typically consists of parallel one- and two-qubit gates, and it is intended to implement a $2^n \times 2^n$ unitary $U(L_i)$. Together, the layers are intended to implement $U(c) = U(L_d) \cdots U(L_1)$.

If quantum circuit c is implemented without error, its output bit string b is a sample from a distribution $P(c)$ whose probabilities are given by $\Pr(b = x) = |\langle x|U(c)|00 \cdots 0\rangle|^2$ where $|00 \cdots 0\rangle = |0\rangle \otimes \cdots \otimes |0\rangle$, $|x\rangle = |x_1\rangle \otimes \cdots \otimes |x_n\rangle$, and x_i is the i -th bit in x . However, when a circuit is executed on a real quantum computer, errors can occur and this means that its output bit string b is a sample from some other distribution $\tilde{P}(c)$. The process of errors corrupting a quantum computation can be modelled as follows. Each logic layer L_i implements the intended unitary superoperator $\mathcal{U}(L_i) : \rho \rightarrow U(L_i)\rho U^\dagger(L_i)$, where ρ is a general n -qubit state, followed by an error channel Λ_i that is a completely positive and trace preserving (CPTP) superoperator [Blume-Kohout et al., 2022]. The imperfect implementation of a circuit c is then simply $\tilde{U}(c) = \prod_{i=1}^d \Lambda_i \circ \mathcal{U}(L_i)$, and the output bit string b is x with probability $\Pr(b = x) = \text{Tr}(|x\rangle\langle x|\tilde{U}(c)|00 \cdots 0\rangle\langle 00 \cdots 0|)$.

2.2 Quantum capability learning

Because quantum computers are error-prone, knowing which quantum circuits a particular quantum computer can execute with low error probability is important. Known as *quantum capability learning* [Proctor et al., 2021a, Hothem et al., 2024b], this task formally involves learning the mapping between a set of quantum circuits $c \in \mathcal{C}$ and some success metric $s(c) \in \mathbb{R}$ quantifying how well c runs on a quantum computer $\mathcal{Q}$. In this work, we consider a large class of circuits known as Clifford (or stabilizer) circuits [Aaronson and Gottesman, 2004], which are sufficient to enable quantum error correction [Campbell et al., 2017], and two widely-used success metrics: *probability of successful trial* (PST) (a.k.a. success probability) and the *process fidelity* (a.k.a. entanglement fidelity) [Hothem et al., 2024b, Nielsen, 2002].

PST is defined only for definite-outcome circuits, which are circuits whose output distribution has (if run without error) support on a single bit string, $b(c)$. For any such circuit c , PST is defined as

$$\text{PST}(c) = \Pr(\text{measuring } b(c) \text{ when executing } c \text{ on } \mathcal{Q}). \quad (1)$$

In practice, $\text{PST}(c)$ is estimated by running the circuit $N_{\text{shots}} \gg 1$ times on $\mathcal{Q}$ and calculating

$$\widehat{\text{PST}(c)} = \frac{\# \text{ observations of } b(c)}{N_{\text{shots}}}. \quad (2)$$

Process fidelity is defined for all circuits, and it quantifies how close the actual quantum evolution of the qubits is to the ideal unitary evolution. It is given by

$$F(c) = \frac{1}{4^n} \text{Tr} [\tilde{U}(c) \mathcal{U}^{-1}(c)]. \quad (3)$$

Estimating $F(c)$ is more complicated than estimating $\text{PST}(c)$, but efficient methods exist, such as mirror circuit fidelity estimation [Proctor et al., 2022]. Hence, in theory, it is possible to efficiently gather training data using either $\text{PST}(c)$ or $F(c)$ on arbitrarily large quantum computers.

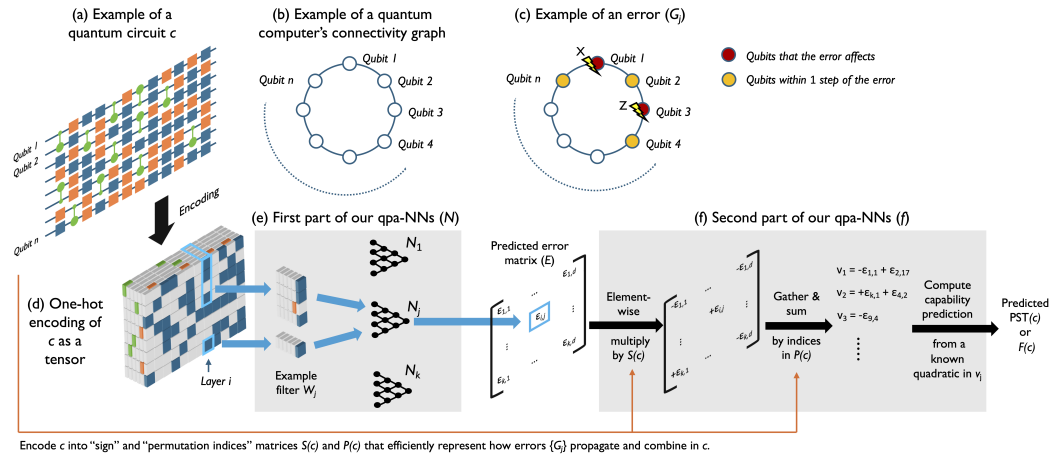


Figure 1: Quantum capability learning with quantum-physics-aware neural networks (qpa-NNs). Our qpa-NNs are a novel architecture for learning a quantum computer's capability, i.e., the mapping from quantum circuits (or programs) to how well that imperfect quantum computer can run those circuits. These networks build in physical principles for how errors in quantum circuits occur—which can be expressed in terms of a quantum computer's connectivity graph—and efficient approximations to the physics of how these errors combine to impact a circuit's success rate.

2.3 Modelling errors in quantum computers

Our qpa-NNs build in efficient approximations to the quantum physics of errors in quantum computers. They do so using the following parameterization of an error channel: $\Lambda = \exp(\sum_j \epsilon_j G_j)$. Here $\mathcal{G}_n = \{G_j\}$ is the set of $2^{2n+1} - 2$ different Hamiltonian (H) and Stochastic (S) *elementary error generators* introduced by Blume-Kohout et al. [2022], and ϵ_j is the *rate* of error G_j . Not every kind of error process can be represented in this form (e.g., amplitude damping, or non-Markovian errors), but this parameterization includes many of the most important kinds of errors in contemporary quantum computers. Each H and S error generator is indexed by a non-identity element of the n -qubit Pauli group ($\mathbb{P}_n$). The Pauli operator indexing an H or S error indicates the qubits it impacts and its direction, e.g., the H error generator indexed by $X \otimes I^{\otimes(n-1)}$ is a coherent error on the 1st qubit and it rotates that qubit around its X axis.

Our qpa-NNs use approximate formulas for computing $\text{PST}(c)$ or $F(c)$ from the rates of H and S errors, which we now review. Consider pushing each error channel Λ_i to the end of the circuit and combining them together, i.e., we compute the error channel $\Lambda(c)$ defined by $\tilde{\mathcal{U}}(c) = \Lambda(c) \circ \mathcal{U}(c)$. Then

$$\text{PST}(c) \approx 1 - \sum_{P \in \mathbb{P}_n^{X,Y}} \left(s_P + \theta_P^2 \right), \quad (4)$$

where s_P and θ_P are the rates of the P -indexed S and H error generators, respectively, in c 's error channel $\Lambda(c)$, and $\mathbb{P}_n^{X,Y}$ is the set of n -qubit Pauli operators containing at least one X or Y . Similarly,

$$F(c) \approx 1 - \sum_{P \in \mathbb{P}_n} \left(s_P + \theta_P^2 \right). \quad (5)$$

Equations (4) and (5) are good approximations for low-error circuits [Mądzik et al., 2022]. However, they both suffer from the same flaw: they require tracking $\mathcal{O}(4^n)$ parameters. To address this problem, our qpa-NNs make an approximation: they only account for the contributions of a polynomially-sized set of errors that contains all those errors which are most likely to be experienced by a quantum computer. In this work, we chose to account for only local, low-weight errors, i.e., those with initial support on a small, connected subset of a device's connectivity graph.

3 Neural network architecture

Our neural network architecture (see Fig. 1) for quantum capability learning combines neural network layers that have GNN-like structures with efficient approximations to the physics of errors in quantum computers. The overall action of our neural networks is to map an encoding of a circuit c to a prediction for $\text{PST}(c)$ or $F(c)$. The same network can predict either $\text{PST}(c)$ or $F(c)$ by simply toggling between two different output layers that have no trainable parameters. Our architecture is divided into two sequential parts. The first part of our architecture is a neural network $\mathcal{N}$ that has the task of learning about the kinds and rates of errors that occur in quantum circuits. We use GNN-like structures within $\mathcal{N}$ to embed physics knowledge for how those errors depend on the quantum circuit being run. The second part of our architecture is a function f with no learnable parameters, that turns $\mathcal{N}$'s output into a prediction for $\text{PST}(c)$ or $F(c)$.

3.1 Physics-aware neural networks for predicting errors in quantum circuits

The neural network $\mathcal{N}$'s input is a quantum circuit c of depth $d(c)$ represented by (i) a tensor $I(c) \in \{0, 1\}^{n \times d(c) \times n_{ch}}$ describing the gates in c (see Fig. 1a), and (ii) a matrix $M(c) \in \{0, 1\}^{2 \times n}$ describing the measurement of the qubits at the end of c . $\mathcal{N}$ maps $I(c)$ to a matrix $\mathcal{E} \in \mathbb{R}^{k \times d(c)}$ and $M(c)$ to a vector $\vec{m} \in \mathbb{R}^k$. $\mathcal{E}_{ij}$ is a prediction for the rate with which error type j occurs during circuit layer i , and m_j is a prediction for the rate with which error type j occurs when measuring the qubits at the end of a circuit. There are $2(4^n - 1)$ different possible error types that can occur in principle (see Section 2) so it is infeasible to predict all their rates beyond very small n . However, the overwhelming majority of these errors are implausible, i.e., they are not expected to occur in real quantum computers [Blume-Kohout et al., 2022]. Our networks therefore predict the rates of every error from a relatively small set of error types $\mathbb{G} = \{G_1, \dots, G_k\}$ containing the k most plausible kinds of error. $\mathbb{G}$ is a hyperparameter of our networks. It can be chosen to reflect the known physics of a particular quantum computer and/or optimized using hyperparameter tuning. In our demonstrations, we choose $\mathbb{G}$ to contain all one-body H and S errors as well as all two-body H and S errors that interact pairs of qubits within h steps on the modelled quantum computer's connectivity graph for some constant h (see Fig. 1b-c, where Fig. 1c shows an H or S error in $\mathbb{G}$ if $h \geq 2$). This choice for $\mathbb{G}$ encodes the physical principles that errors are primarily either localized to a qubit or are two-body interactions between nearby qubits [Blume-Kohout et al., 2022]. The size of $\mathbb{G}$ grows with n , and for planar connectivity graphs (as in, e.g., contemporary superconducting qubit systems [Arute et al., 2019]) it grows linearly in n . This results in $k = \mathcal{O}(n)$ errors whose rates $\mathcal{N}$ must learn to predict.

The internal structures of $\mathcal{N}$ are chosen to reflect general physical principles for how $\mathcal{E}$ and $\vec{m}$ depend on c . $\mathcal{E}_{ij}$ is a prediction for the rate that G_j occurs in circuit layer i , and this error corresponds to a space/time location within c —because it occurs at layer index or time i and G_j acts on a subset of the qubits $Q(G_j)$ (see example in Fig. 1c). This error's rate will therefore primarily depend only on the gates in a time- and space-local region around its location in c . Furthermore, this dependence will typically be invariant under time translations (this is true except for some exotic non-Markovian kinds of errors, which we discuss in Section 8.1). We can encode these structures into $\mathcal{N}$ by predicting $\mathcal{E}_{ij}$ from a space-time “window” of c around the associated error's location using a filter W_j that “slides” across the circuit to predict the rate of G_j versus time i . Stated more formally, we predict $\mathcal{E}_{ij}$ using a multilayer perceptron $\mathcal{N}_j$ whereby $\mathcal{N}_j(W_j[I(c), i]) = \mathcal{E}_{ij}$ and $W_j[I(c), i]$ is a snippet of $I(c)$ whose temporal origin is i (see Fig. 1e). The shape of each filter W_j is a hyperparameter of our networks and it can be designed to reflect general physical principles, the known physics of a particular quantum computing system, and/or optimized with hyperparameter tuning. The particular neural networks we present later herein use filters $W_j(I(c), i)$ that snip out only layer i and discard the parts of the layer that act on qubits more than l steps away from $Q(G_j)$ in the quantum computer's connectivity graph (e.g., the filter shown in Fig. 1e corresponds to the error shown in Fig. 1c and $l = 1$). This neural network structure has close connections to graph convolution layers [Kipf and Welling, 2016], as well as CNNs. We choose this structure as it can model spatially localized crosstalk errors, which are a ubiquitous but hard-to-model class of errors in quantum computers [Sarovar et al., 2020].

The network $\mathcal{N}$ must also predict the rates of errors that occur during measurements (unless the qpa-NN will only ever predict $F(c)$ not $\text{PST}(c)$), but these are typically independent of the rates of gate errors (which are predicted by the $\mathcal{N}_j$). So we do not use the $\mathcal{N}_j$ and their convolutional filters W_j to make predictions for $\vec{m}$. Instead we use separate but structurally equivalent networks $\mathcal{N}'_j$ with

corresponding filters W'_j that take $M(c)$ as input and implement only spatial filtering. That is, W'_j simply discards rows from $M(c)$, as, unlike $I(c)$, $M(c)$ has no temporal dimension. The W'_j are hyperparameters of our networks allowing us to separately adjust the shape of each W'_j to reflect the known physics of errors induced by measuring qubits. In our demonstrations, our W'_j filters have the same structure as the W_j filters but with an independent l' steps parameter (large l' enables modelling many-qubit measurement crosstalk).

3.2 Processing predicted error rates to predict capabilities

We process $\mathcal{N}$'s output to predict $\text{PST}(c)$ or $F(c)$ using a function f with no learnable parameters. This turns $\mathcal{N}$'s output into the two quantities of interest, and it also makes training $\mathcal{N}$ feasible. We cannot easily train $\mathcal{N}$ in isolation because the error matrix $\mathcal{E}$ predicted by $\mathcal{N}$ is not a directly observable quantity. Generating the data needed to train $\mathcal{N}$ directly would require extraordinarily expensive quantum process tomography [Nielsen et al., 2021], which is infeasible except for very small n . In contrast, both $\text{PST}(c)$ and $F(c)$ can be efficiently estimated (see Section 2) for a given circuit c .

The function f computes an approximation to the value for $\text{PST}(c)$ or $F(c)$ predicted by $\mathcal{E}$ and $\vec{m}$. The matrix $\mathcal{E}$ encodes the prediction that c 's imperfect action is

$$\tilde{U}(c) = \Lambda_d(\mathcal{E})\mathcal{U}(L_d) \cdots \Lambda_1(\mathcal{E})\mathcal{U}(L_1), \quad (6)$$

where the L_i are the d layers of c (see Section 2) and $\Lambda_i(\mathcal{E}) = \exp(\sum_{j=1}^k \mathcal{E}_{ij} G_j)$, i.e., $\Lambda_i(\mathcal{E})$ is an error channel parameterized by the i^{th} column of $\mathcal{E}$. Equation (6) implies an exact prediction for $\text{PST}(c)$ or $F(c)$ [e.g., Eq. (3)], but exactly computing that prediction involves explicitly creating and multiplying together each of the $4^n \times 4^n$ matrices in Eq. (6). This is infeasible, except for very small n . Instead our f computes an efficient approximation to this prediction.

Our function f 's action is most easily described by embedding $\mathcal{E}$ into the space of all possible H and S errors $\mathbb{G}_n$, resulting in a $d \times (2^{2n+1} - 2)$ matrix $\mathcal{E}_e$ whose columns are k -sparse. However, we never construct these exponentially large matrices. Consider pulling each error channel to the end of the circuit, giving $\tilde{U}(c) = \Lambda'_d(\mathcal{E}'_e) \cdots \Lambda'_1(\mathcal{E}'_e)\mathcal{U}(c)$ where $\Lambda'_d(\mathcal{E}'_e) = \exp(\sum_{j=1}^{2^{2n+1}-2} [\mathcal{E}'_e]_{dj} G_j)$. Because c contains only Clifford gates and Clifford unitaries preserve the Pauli group [Aaronson and Gottesman, 2004], $\mathcal{E}'_e$ has columns that are just c -dependent signed permutations of $\mathcal{E}_e$'s columns. The signed permutations required can be efficiently computed in advance (i.e., as an input encoding step) using an efficient representation of Clifford unitaries [Gidney, 2021]. Furthermore, these permutations can be efficiently represented in two $d \times k$ matrices: a *sign matrix* $S(c)$ containing ± 1 signs to be element-wise multiplied with $\mathcal{E}$ and a *permutation indices matrix* $P(c)$ containing integers between 1 and $2^{2n+1} - 2$, where P_{ij} specifies what error G_j becomes when pulled through the $d - i$ circuit layers after layer i .

We now have a representation of $\mathcal{E}$'s prediction for the circuit c 's error map $\Lambda(c)$ as a sequence of error maps $\Lambda'_d(\mathcal{E}'_e) \cdots \Lambda'_1(\mathcal{E}'_e)$, and we need to predict $\text{PST}(c)$ or $F(c)$. We can do so if we can compute $\mathcal{E}$'s prediction for the S and H error rates in $\Lambda(c)$, as we can then apply Eq. (4) or Eq. (5). To achieve this, we combine the $\Lambda'_i(\mathcal{E})$ into a single error map using a first-order Baker-Campbell-Hausdorff (BCH) expansion. Using our embedded representation, this means simply approximating $\Lambda(c)$ as $\Lambda(c) \approx \exp(\sum_j v_j G'_j)$ where $v_j = \sum_{i=1}^d [\mathcal{E}'_e]_{ij}$, i.e., we sum over the rows of $\mathcal{E}'_e$. To predict $F(c)$ we then simply apply Eq. (4) (meaning summing up v_j with those elements that correspond to Hamiltonian errors squared). Because measurement errors impact $\text{PST}(c)$, to predict $\text{PST}(c)$ we again apply the BCH expansion to combine in the predicted measurement error map $\exp(\sum_{j=1}^l m_j G_j)$ and then apply Eq. (5). The efficient representation of the overall action of f is illustrated in Fig. 1 (the addition of the measurement error map is not shown).

4 Datasets

4.1 Experimental 5-qubit data

We used the 5-qubit datasets from Hothem et al. [2024b] for our experimental demonstrations. Each of these datasets $D = \{(c, \widehat{\text{PST}(c)})\}$ was gathered by running random and periodic mirror circuits

(two types of definite-outcome circuits) on 5-qubit IBM Q computers (ibmq_london, ibmq_essex, ibmq_burlington, ibmq_vigo, ibmq_ourense and ibmq_yorktown), and estimating the PST of each circuit using Eq. (2). Each circuit was run between 1024 and 4096 times, with the exact number depending upon how many times the circuit sampling process generated the circuit (some short, 1-qubit circuits were generated multiple times). The random and periodic mirror circuits contained between 1 and 5 active qubits—called the circuit’s *width*—and ranged in depth from 3 to 515 layers (alt. 259 layers for the ibmq_yorktown dataset).

As we focus on high-PST circuits, we removed all circuits with a PST less than 85% from each dataset, leaving between 864 (ibmq_burlington) and 1369 (ibmq_yorktown) circuits in each dataset. The remaining circuits were partitioned into training, validation, and test sets by their original assignment in Hothem et al. [2024b]. This setup enables a direct comparison between our qpa-NNs and the CNNs trained in Hothem et al. [2024b]. Training set sizes ranged from 682 circuits on ibmq_burlington to 1097 circuits on ibmq_yorktown, with an approximate training, validation, testing split of 80%, 10%, and 10%, respectively.

4.2 Simulated 4-qubit data

For our 4-qubit simulations, we generated 5 datasets of 5000 high-fidelity ($F(c) > 85\%$) random circuits, for a hypothetical 4-qubit processor with a “ring” geometry (i.e., like that in Fig. 1b). The circuits ranged in width (w) from 1 to 4 qubits, and in depth from 1 to 180 circuit layers. We designed each circuit for a randomly chosen subset of w qubits. Each circuit layer was created by *i.i.d.* sampling from all possible circuit layers on the w active qubits. We used a gate set containing two-qubit CNOT gates and 7 different single-qubit gates (specifically $\{X(\pi/2), Y(\pi/2), X(3\pi/2), Y(3\pi/2), X(\pi), Y(\pi), Z(\pi)\}$ where $P(\theta)$ denotes a rotation around the P axis of the Bloch sphere by θ). See Appendix C for additional details.

All circuits were simulated under the same error model, consisting of local coherent (i.e., H) errors, to exactly compute each c ’s $F(c)$ [Fig. 3 shows a histogram of $F(c)$]. After removing duplicate circuits, the resulting datasets $D = \{(c, F(c))\}$ were partitioned into training, validation, and testing subsets, with a partition of 56.25%, 18.75%, and 25%, respectively. The parameters of the error model were randomly selected: each gate was assigned a small error strength, which was then distributed randomly across all possible (local) one- or two-qubit coherent errors, for the one- and two-qubit gates, respectively. We chose a model with only coherent errors as these errors are ubiquitous, they are hard to model accurately and efficiently, and we conjecture that qpa-NNs can model them.

We also generated 5 datasets of 750 random mirror circuits on the same hypothetical 4-qubit quantum computer. Again, the random mirror circuits varied in width from 1 to 4 qubits, and were designed to be run on a randomly selected subset of w qubits. However, instead of *i.i.d.* sampling of each circuit layer, each circuit was randomly sampled from the class of random mirror circuits on the w qubits. The depth of the mirror circuits ranged from 8 to 174 layers. Because we generated the mirror circuit datasets to evaluate how well qpa-NNs and CNNs generalize to out-of-distribution circuits, they were used exclusively as testing sets. To ensure that no training was performed on mirror circuits, we removed any mirror circuits that appeared in the random circuit sets (in actuality, there were no duplicates).

4.3 Simulated 100-qubit data

For our 100-qubit simulation, we generated a single dataset of 5000 high-fidelity ($F(c) > 91\%$) random circuits, for a hypothetical 100-qubit quantum computer with a “ring” geometry. All of the circuits had a width of 100 qubits, and ranged in depth from 1 to 22 circuit layers. We sampled circuit layers using the same process and gate set as in the 4-qubit simulations.

We simulated every circuit using the same error model, consisting of local, weight-1 S and H errors. As before, the parameters of the error model were randomly selected and the data were partitioned into training, validation, and testing subsets according to a 56.25%, 18.75%, 25% split.

Unlike in our 4-qubit simulations, we did not compute $F(c)$ exactly as doing so for a 100-qubit circuit is infeasible in the presence of coherent errors. Instead, we used a first-order simulation method to approximate $F(c)$. In this method, $F(c)$ is computed by assigning each gate its own error vector based on the error model, adding up the error vectors layer-wise to compute an error vector

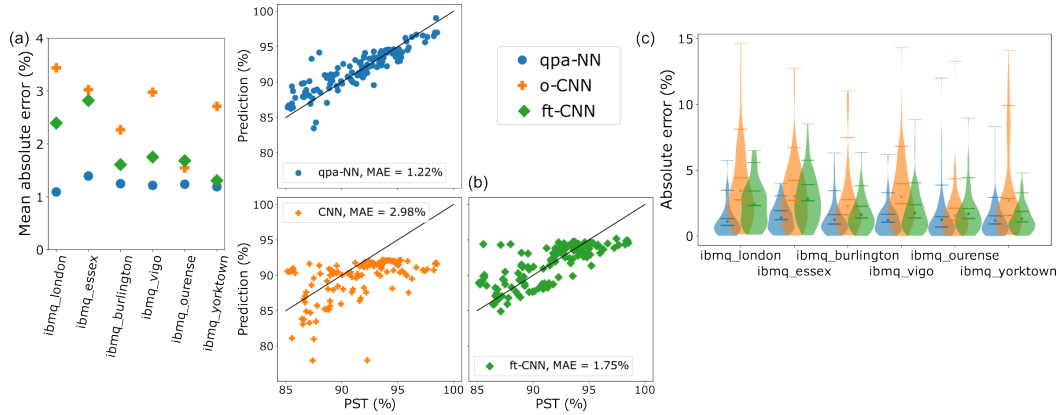


Figure 2: **Prediction accuracy on real quantum computers.** (a) The mean absolute error of our qpa-NNs (●), the CNNs from Hothem et al. [2024b] (o-CNN, +), and fine-tuned CNNs (ft-CNN, ♦) on the test data. (b) The predictions of the three models for `ibmq_vigo` on the test data, and (c) the distribution of each model’s absolute error on the test data, including the 50th, 75th, 95th and 100th percentiles (lines) and the means (points).

for each circuit layer, and then computing $F(c)$ as in the second part of a qpa-NN [Figure 1(f)]. See Appendix C.4 for more details.

4.4 Encoding schemes

We used two different encoding schemes for converting each circuit c into a tensor. For the CNNs on experimental data, we used the same encoding scheme as Hothem et al. [2024b], as we used their data and networks. For all qpa-NNs, and the CNNs on simulated data, we used the following scheme. As outlined in Section 3, each width- w circuit c is represented by a three-dimensional tensor $I(c) \in \{0, 1\}^{n \times d(c) \times n_{ch}}$ describing the gates in c and a matrix $M(C) \in \{0, 1\}^{2 \times w}$ describing the measurement of the qubits. The ij -th entry of $I(c)$,

$$I_{ij}(c) = (I_{ij1}(c), \dots, I_{ijn_{ch}}(c)), \quad (7)$$

is a one-hot encoded vector of what happens to qubit i in layer j . For the hypothetical 4-qubit ring processor, $n_{ch} = 11$: one channel for each single-qubit gate and four channels for the CNOT gates. There are four CNOT channels to specify if the qubit i is the target or control qubit and if the interacting qubit is to the left or right of qubit i . We used an additional 4 or 8 CNOT channels for the experimental data, depending on the quantum computer’s geometry. The first row in $M(c)$ is the bitstring specifying which qubits are measured at the end of c . When c is a definite-outcome circuit, the second row is its target bit string, i.e., the sole bit string in the support of c ’s outcome distribution when it is executed without error [i.e., $P(c)$]. Both $I(c)$ and $M(c)$ are zero-padded to ensure a consistent tensor shape across a dataset.

Additionally, each circuit c is accompanied by a permutation matrix $P(c) \in \mathbb{N}^{n \times k}$ and sign matrix $S(c) \in \{\pm 1\}^{n \times k}$. The ij -entry of $P(c)$ specifies which error the j -th tracked error occurring after the i -th layer is transformed into at the end of the circuit. The ij -th entry of $S(c)$ specifies the sign of that error.

5 5-qubit experiments

We now present the results from our head-to-head comparison between the qpa-NNs and the CNNs on the 5-qubit datasets used in Hothem et al. [2024b]. Figure 2 shows the mean absolute error (MAE) achieved by the CNNs (+) and the qpa-NNs (●) on each of the datasets. For all datasets, MAE is lower for the qpa-NNs than the CNNs, with an average reduction of 50.4% ($\sigma_x = 16.7\%$, i.e., the standard deviation of percent-drop in MAE). The Bayes factor K is between $K = 10^{30}$ and $K = 10^{383}$ (here, K is the ratio of the likelihood of the qpa-NN to the likelihood of the CNN given the test data). This is overwhelming evidence that the qpa-NN is a better model ($K \geq 10^2$ is typically

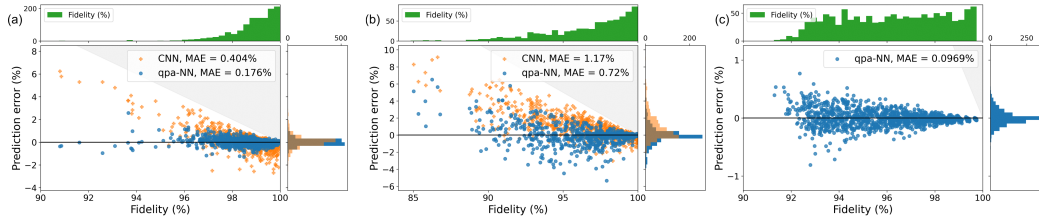


Figure 3: Demonstrating our qpa-NNs’ accuracy for hard-to-model coherent errors and at scale. (a) Scatter plot of the prediction errors on test data of a qpa-NN (●) and CNN (●) trained to predict the fidelity $F(c)$ of random circuits run on a hypothetical 4-qubit quantum computer. The qpa-NN significantly outperforms the CNN. The top subplot contains a histogram (green bars) of the ground-truth fidelities. (b) Prediction errors on out-of-distribution test data, from random mirror circuits. The qpa-NN achieves modest prediction accuracy on this out-of-distribution task, suggesting that the qpa-NNs are accurately learning error rates. (c) Prediction errors on the 100-qubit test data, demonstrating that our qpa-NN approach can accurately predict $F(c)$ for circuits run on large-scale quantum computers.

considered decisive). These results strongly suggest that the extra infrastructure in the qpa-NNs is making a difference.

The improved performance of the qpa-NNs is not because of an increase in model size. For example, the `ibmq_london` CNN contains 6,649,531 trainable parameters compared to the 1,218,348 trainable parameters in the qpa-NN. Moreover, CNNs of similar or larger sizes than the qpa-NNs were included in the hyperparameter optimization space of the CNNs [Hothem et al., 2024b].

Nonetheless, comparing the qpa-NNs to the CNNs is somewhat unfair as the CNNs were trained on out-of-distribution circuits—they were trained on the entire training dataset from Hothem et al. [2024b] which also contains low-PST circuits. For a fairer comparison, we fine-tuned each CNN (Fig. 2, ♦) on the same high-PST training set used to train the qpa-NNs. Fine-tuning typically increased the CNNs’ performances (mean 25.1% improvement, $\sigma_x = 22.3\%$). However, the qpa-NNs achieve a MAE that is lower than the fine-tuned CNNs by 32.2% on average ($\sigma_x = 17.3\%$) and outperform the fine-tuned CNNs on all six datasets. K is between 10^{28} and 10^{238} , which is overwhelming evidence that the qpa-NNs are better models than the fine-tuned CNNs.

6 4-qubit simulations

One reason why the extra infrastructure in our qpa-NNs may be necessary is that off-the-shelf networks struggle with modeling coherent errors [Hothem et al., 2024b]. To test our hypothesis, we trained a qpa-NN to predict the fidelity $F(c)$ of random circuits executed on a hypothetical 4-qubit quantum computer experiencing purely coherent errors. We compared this qpa-NN to a hyperparameter-tuned CNN trained on the same data. Figure 3 shows the results from one representative dataset.

The qpa-NNs again significantly outperform the CNNs. Across the five datasets, the qpa-NNs’ averaged a 52.4% reduction in MAE ($\sigma_x = 3.00\%$) on the test data. We also see a significant improvement in the mean Pearson correlation coefficient, $\bar{r}_{\text{qpa-NN}} = .968$ vs. $\bar{r}_{\text{CNN}} = .749$.

We also found that qpa-NNs trained on random circuits are modest predictors of the infidelity of random mirror circuits, which are a different family of circuits. This is an example of out-of-distribution generalization. Random mirror circuits differ in a variety of ways from the random circuits on which the qpa-NNs were trained, including both the presence of idle gates (which are noiseless in our simulations) and a motion-reversal structure in the circuits that causes the addition or cancellation of errors that are far apart in time. The qpa-NNs achieve an average MAE of .72% on the random mirror circuits ($\sigma_x = .046\%$). Although this is a $3.2\times$ increase in MAE over the in-distribution test data, the strong linear relation between the network’s predictions and the ground truth ($\bar{r} = .912$, $\sigma_x = .009$) strongly suggests that the qpa-NNs are learning information relevant to random mirror circuits.

7 100-qubit simulation

Quantum-physics-aware neural networks scale just as well as CNNs, despite their extra infrastructure. To demonstrate their scalability, we trained a qpa-NN to predict the fidelity $F(c)$ of random circuits executed on a hypothetical 100-qubit quantum computer experiencing a mix of stochastic and coherent errors. To our knowledge, this is the first creation of a capability model of any kind, for a 100+ qubit quantum computer. Figure 3(c) shows the results from our demonstration.

The qpa-NN achieved a MAE of 0.097%. While the underlying noise model was quite simple, this result shows that it is technically feasible to construct qpa-NN capability models for today's moderate-scale quantum computers and for tomorrow's early fault-tolerant quantum computers.

8 Discussion

8.1 Limitations

Our results are a significant improvement over the state of the art, but our approach does have several limitations:

1. As presently conceived, our approach assumes that the modelled quantum computer's error rates are invariant under time translations, which is a kind of Markovianity assumption (although it is weaker than the typical Markovian assumption used in conventional quantum computer models [Nielsen et al., 2021]). However, non-Markovian noise exists in quantum computers [White et al., 2020]. In the future, we plan to address this issue by adding temporal information into our approach, perhaps with a temporal or positional encoding [Vaswani et al., 2017].
2. Our approach only considers two error classes (H and S errors). Other Markovian error classes, like amplitude damping, exist, but their error rates ε contribute to PST and fidelity at order $\mathcal{O}(\varepsilon^3)$ [Mađzik et al., 2022]. Our approach can be easily extended to include those errors, if necessary, by learning their rates with $\mathcal{N}$ and updating f to account for their presence.
3. Our current approach works for Clifford circuits, which includes arguably the most important kinds of circuits (e.g., quantum error correction circuits) but not all interesting circuits. This is because our method for efficiently propagating errors through circuits (implemented by f together with the S and P matrices) leverages the elegant mathematics of Clifford circuits. Our approach can be easily extended to generic few-qubit quantum circuits ($\lesssim 10$ qubits), but to obtain the efficiency needed for large n with general circuits we will need to develop approximate methods for propagating errors through those circuits.

8.2 Conclusion

In this paper, we presented a new quantum-physics-aware neural network architecture for modelling a quantum computer's capability that significantly improves upon the state of the art. The new architecture concatenates two parts: (i) a neural network with structural similarities to GNNs that uses gate information and a quantum computer's connectivity graph to predict the rates of errors in each of a circuit's layers, and (ii) a non-trainable function that turns the predicted error rates into a capability prediction. By imbuing these networks with knowledge about how errors occur and combine within a circuit, we are able to outperform state-of-the-art CNN-based capability models by $\sim 50\%$ on both experimental data and simulated data. We also provided evidence that our quantum-physics-aware networks are learning the true physical error rates, as they exhibit modest prediction accuracy when predicting the fidelity of out-of-distribution quantum circuits, which would enable our networks to also be used to diagnose the error processes occurring in a particular quantum computer (an important task known as characterization or tomography [Nielsen et al., 2021]).

Understanding which quantum circuits a quantum computer can run, and how well it can run them, is an important yet challenging component of understanding a quantum computer's power. Given the complexity of the problem, neural networks are likely to play a large role in its solution. As our results demonstrate, our new physics-aware network architecture could play a critical role in building fast and reliable neural network-based capability models.

Acknowledgments and Disclosure of Funding

This material was funded in part by the U.S. Department of Energy, Office of Science, Office of Advanced Scientific Computing Research, Quantum Testbed Pathfinder Program, and by the Laboratory Directed Research and Development program at Sandia National Laboratories. T.P. acknowledges support from an Office of Advanced Scientific Computing Research Early Career Award. We acknowledge the use of IBM Quantum services for this work. The views expressed are those of the authors, and do not reflect the official policy or position of IBM or the IBM Quantum team.

Sandia National Laboratories is a multi-mission laboratory managed and operated by National Technology & Engineering Solutions of Sandia, LLC (NTESS), a wholly owned subsidiary of Honeywell International Inc., for the U.S. Department of Energy's National Nuclear Security Administration (DOE/NNSA) under contract DE-NA0003525. This written work is authored by an employee of NTESS. The employee, not NTESS, owns the right, title and interest in and to the written work and is responsible for its contents. Any subjective views or opinions that might be expressed in the written work do not necessarily represent the views of the U.S. Government. The publisher acknowledges that the U.S. Government retains a non-exclusive, paid-up, irrevocable, world-wide license to publish or reproduce the published form of this written work or allow others to do so, for U.S. Government purposes. The DOE will provide public access to results of federally sponsored research in accordance with the DOE Public Access Plan.

References

- S. Aaronson and D. Gottesman. Improved simulation of stabilizer circuits. *Phys. Rev. A*, 70(5):052328, Nov. 2004. ISSN 1050-2947. doi: 10.1103/PhysRevA.70.052328. URL <https://link.aps.org/doi/10.1103/PhysRevA.70.052328>.
- M. Abadi, A. Agarwal, P. Barham, E. Brevdo, Z. Chen, C. Citro, G. S. Corrado, A. Davis, J. Dean, M. Devin, S. Ghemawat, I. Goodfellow, A. Harp, G. Irving, M. Isard, Y. Jia, R. Jozefowicz, L. Kaiser, M. Kudlur, J. Levenberg, D. Mané, R. Monga, S. Moore, D. Murray, C. Olah, M. Schuster, J. Shlens, B. Steiner, I. Sutskever, K. Talwar, P. Tucker, V. Vanhoucke, V. Vasudevan, F. Viégas, O. Vinyals, P. Warden, M. Wattenberg, M. Wicke, Y. Yu, and X. Zheng. TensorFlow: Large-scale machine learning on heterogeneous systems, 2015. URL <http://tensorflow.org/>. Software available from tensorflow.org.
- F. Arute, K. Arya, R. Babbush, D. Bacon, J. C. Bardin, R. Barends, R. Biswas, S. Boixo, F. G. S. L. Brandao, D. A. Buell, B. Burkett, Y. Chen, Z. Chen, B. Chiaro, R. Collins, W. Courtney, A. Dunsworth, E. Farhi, B. Foxen, A. Fowler, C. Gidney, M. Giustina, R. Graff, K. Guerin, S. Habegger, M. P. Harrigan, M. J. Hartmann, A. Ho, M. Hoffmann, T. Huang, T. S. Humble, S. V. Isakov, E. Jeffrey, Z. Jiang, D. Kafri, K. Kechedzhi, J. Kelly, P. V. Klimov, S. Knysh, A. Korotkov, F. Kostritsa, D. Landhuis, M. Lindmark, E. Lucero, D. Lyakh, S. Mandrà, J. R. McClean, M. McEwen, A. Megrant, X. Mi, K. Michielsen, M. Mohseni, J. Mutus, O. Naaman, M. Neeley, C. Neill, M. Y. Niu, E. Ostby, A. Petukhov, J. C. Platt, C. Quintana, E. G. Rieffel, P. Roushan, N. C. Rubin, D. Sank, K. J. Satzinger, V. Smelyanskiy, K. J. Sung, M. D. Trevithick, A. Vainsencher, B. Villalonga, T. White, Z. J. Yao, P. Yeh, A. Zalcman, H. Neven, and J. M. Martinis. Quantum supremacy using a programmable superconducting processor. *Nature*, 574(7779):505–510, Oct. 2019. ISSN 0028-0836, 1476-4687. doi: 10.1038/s41586-019-1666-5. URL <http://dx.doi.org/10.1038/s41586-019-1666-5>.
- R. Blume-Kohout, M. P. da Silva, E. Nielsen, T. Proctor, K. Rudinger, M. Sarovar, and K. Young. A taxonomy of small markovian errors. *PRX Quantum*, 3:020335, May 2022. doi: 10.1103/PRXQuantum.3.020335. URL <https://link.aps.org/doi/10.1103/PRXQuantum.3.020335>.
- E. T. Campbell, B. M. Terhal, and C. Vuillot. Roads towards fault-tolerant universal quantum computation. *Nature*, 549(7671):172–179, Sept. 2017. ISSN 0028-0836, 1476-4687. doi: 10.1038/nature23460. URL <http://dx.doi.org/10.1038/nature23460>.
- Y. Cao, J. Romero, J. P. Olson, M. Degroote, P. D. Johnson, M. Kieferová, I. D. Kivlichan, T. Menke, B. Peropadre, N. P. D. Sawaya, S. Sim, L. Veis, and A. Aspuru-Guzik. Quantum chemistry in the age

- of quantum computing. *Chemical Reviews*, 119(19):10856–10915, 2019. doi: 10.1021/acs.chemrev.8b00803. URL <https://doi.org/10.1021/acs.chemrev.8b00803>. PMID: 31469277.
- F. Chollet et al. Keras. <https://keras.io>, 2015.
- N. Elsayed Amer, W. Gomaa, K. Kimura, K. Ueda, and A. El-Mahdy. On the learnability of quantum state fidelity. *EPJ Quantum Technology*, 9:31, 2022. URL <https://epjquantumtechnology.springeropen.com/articles/10.1140/epjqt/s40507-022-00149-8#citeas>.
- C. Gidney. Stim: a fast stabilizer circuit simulator. *Quantum*, 5:497, July 2021. ISSN 2521-327X. doi: 10.22331/q-2021-07-06-497. URL <https://doi.org/10.22331/q-2021-07-06-497>.
- A. W. Harrow, A. Hassidim, and S. Lloyd. Quantum algorithm for linear systems of equations. *Phys. Rev. Lett.*, 103:150502, Oct 2009. doi: 10.1103/PhysRevLett.103.150502. URL <https://link.aps.org/doi/10.1103/PhysRevLett.103.150502>.
- A. Hashim, L. B. Nguyen, N. Goss, B. Marinelli, R. K. Naik, T. Chistolini, J. Hines, J. P. Marceaux, Y. Kim, P. Gokhale, T. Tomesh, S. Chen, L. Jiang, S. Ferracin, K. Rudinger, T. Proctor, K. C. Young, R. Blume-Kohout, and I. Siddiqi. A practical introduction to benchmarking and characterization of quantum computers. *arXiv [quant-ph]*, Aug. 2024. URL <https://arxiv.org/abs/2408.12064>.
- D. Hothem, T. Catanach, K. Young, and T. Proctor. Learning a quantum computer’s capability using convolutional neural networks [Data set]. <https://doi.org/10.5281/zenodo.7829489>, 2023a. Published: 2023-04-12.
- D. Hothem, J. Hines, K. Nataraj, R. Blume-Kohout, and T. Proctor. Predictive models from quantum computer benchmarks. In *2023 IEEE International Conference on Quantum Computing and Engineering (QCE)*, pages 709–714, Los Alamitos, CA, USA, sep 2023b. IEEE Computer Society. doi: 10.1109/QCE57702.2023.00086. URL <https://doi.ieeecomputersociety.org/10.1109/QCE57702.2023.00086>.
- D. Hothem, A. Miller, and T. Proctor. Supplementary material for What is my quantum computer good for? Quantum capability learning with physics-aware neural networks. Online, 2024a. To be released as part of praqtice, a GitHub repository.
- D. Hothem, K. Young, T. Catanach, and T. Proctor. Learning a Quantum Computer’s Capability. *IEEE Transactions on Quantum Engineering*, 5(01):1–26, Jan. 2024b. ISSN 2689-1808. doi: 10.1109/TQE.2024.3430215. URL <https://doi.ieeecomputersociety.org/10.1109/TQE.2024.3430215>.
- D. Kingma and J. Ba. Adam: A method for stochastic optimization. In Y. Bengio and Y. LeCun, editors, *Proceedings of the 3rd International Conference for Learning Representations*. Microtome, 2015. URL <https://arxiv.org/abs/1412.6980>.
- T. N. Kipf and M. Welling. Semi-Supervised classification with graph convolutional networks. *Proceedings of the 5th International Conference for Learning Representations*, Sept. 2016. URL <http://arxiv.org/abs/1609.02907>.
- T. Lubinski, S. Johri, P. Varosy, J. Coleman, L. Zhao, J. Necaie, C. H. Baldwin, K. Mayer, and T. Proctor. Application-Oriented performance benchmarks for quantum computing. *IEEE Transactions on Quantum Engineering*, 4:1–32, 2023. ISSN 2689-1808. doi: 10.1109/TQE.2023.3253761. URL <http://dx.doi.org/10.1109/TQE.2023.3253761>.
- M. T. Mądzik, S. Asaad, A. Youssry, B. Joecker, K. M. Rudinger, E. Nielsen, K. C. Young, T. J. Proctor, A. D. Baczewski, A. Laucht, V. Schmitt, F. E. Hudson, K. M. Itoh, A. M. Jakob, B. C. Johnson, D. N. Jamieson, A. S. Dzurak, C. Ferrie, R. Blume-Kohout, and A. Morello. Precision tomography of a three-qubit donor quantum processor in silicon. *Nature*, 601(7893):348–353, Jan. 2022.
- E. Nielsen, K. Rudinger, T. Proctor, A. Russo, K. Young, and R. Blume-Kohout. Probing quantum processor performance with pyGSTi. *Quantum Sci. Technol.*, 5(4):044002, July 2020. ISSN 2058-9565. doi: 10.1088/2058-9565/ab8aa4. URL <https://iopscience.iop.org/article/10.1088/2058-9565/ab8aa4>.

- E. Nielsen, J. K. Gamble, K. Rudinger, T. Scholten, K. Young, and R. Blume-Kohout. Gate set tomography. *Quantum*, 5(557):557, Oct. 2021. ISSN 2521-327X. doi: 10.22331/q-2021-10-05-557. URL <https://quantum-journal.org/papers/q-2021-10-05-557/>.
- M. A. Nielsen. A simple formula for the average gate fidelity of a quantum dynamical operation. *Physics Letters A*, 303(4):249–252, 2002.
- M. A. Nielsen and I. L. Chuang. *Quantum Computation and Quantum Information: 10th Anniversary Edition*. Cambridge University Press, 2010.
- T. Proctor, K. Rudinger, K. Young, E. Nielsen, and R. Blume-Kohout. Measuring the capabilities of quantum computers. *Nature Phys*, 18(1):75, Dec. 2021a. ISSN 1745-2473. doi: 10.1038/s41567-021-01409-7. URL <https://www.nature.com/articles/s41567-021-01409-7>.
- T. Proctor, K. Rudinger, K. Young, E. Nielsen, and R. Blume-Kohout. Scalable randomized benchmarking of quantum computers using mirror circuits [Data set]. <https://doi.org/10.5281/zenodo.5197499>, 2021b. Accessed: 2023-04-12.
- T. Proctor, S. Seritan, E. Nielsen, K. Rudinger, K. Young, R. Blume-Kohout, and M. Sarovar. Establishing trust in quantum computations, 2022. URL <https://arxiv.org/abs/2204.07568>.
- T. Proctor, K. Young, A. D. Baczewski, and R. Blume-Kohout. Benchmarking quantum computers. *arXiv [quant-ph]*, July 2024. URL <https://arxiv.org/abs/2407.08828>.
- N. C. Rubin, D. W. Berry, A. Kononov, F. D. Malone, T. Khattar, A. White, J. Lee, H. Neven, R. Babbush, and A. D. Baczewski. Quantum computation of stopping power for inertial fusion target design. *Proceedings of the National Academy of Sciences*, 121(23):e2317772121, 2024. doi: 10.1073/pnas.2317772121. URL <https://www.pnas.org/doi/abs/10.1073/pnas.2317772121>.
- M. Sarovar, T. Proctor, K. Rudinger, K. Young, E. Nielsen, and R. Blume-Kohout. Detecting crosstalk errors in quantum information processors. *Quantum*, 4(321):321, Sept. 2020. ISSN 2521-327X. doi: 10.22331/q-2020-09-11-321. URL <https://quantum-journal.org/papers/q-2020-09-11-321/>.
- P. W. Shor. Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer. *SIAM Journal on Computing*, 26(5):1484–1509, 1997. doi: 10.1137/S0097539795293172. URL <https://doi.org/10.1137/S0097539795293172>.
- A. Vadali, R. Kshirsagar, P. Shyamsundar, and G. N. Perdue. Quantum circuit fidelity estimation using machine learning. *Quantum Mach. Intell.*, 6(1), June 2024.
- A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin. Attention is all you need. In *Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS’17*, page 6000–6010, Red Hook, NY, USA, 2017. Curran Associates Inc. ISBN 9781510860964.
- H. Wang, Z. Liang, J. Gu, Z. Li, Y. Ding, W. Jiang, Y. Shi, D. Z. Pan, F. T. Chong, and S. Han. TorchQuantum Case Study for Robust Quantum Circuits. In *Proceedings of the 41st IEEE/ACM International Conference on Computer-Aided Design*, volume 30 of ICCAD ’22, page 1–9. ACM, Oct. 2022. doi: 10.1145/3508352.3561118. URL <http://dx.doi.org/10.1145/3508352.3561118>.
- G. A. L. White, C. D. Hill, F. A. Pollock, L. C. L. Hollenberg, and K. Modi. Demonstration of non-markovian process characterisation and control on a quantum processor. *Nature Communications*, 11(1), Dec. 2020. ISSN 2041-1723. doi: 10.1038/s41467-020-20113-3. URL <http://dx.doi.org/10.1038/s41467-020-20113-3>.

A Compute resources

All of the quantum-physics-aware neural networks used in our 5-qubit experiments and 4-qubit simulations were trained using a 6-Core Intel Core i9 processor on a MacBookPro 15.1 with 32GB of memory. Each model took roughly 15-20 wall clock minutes to train. Total training time, across the paper, totaled ~ 160 wall clock minutes.

All of the 4-qubit simulations and data pre-processing were performed using a 6-core Intel Core i9 processor on a MacBookPro 15.1 with 32GB of memory. Each dataset took approximately 1 hour of wall clock time to create. This total includes the initial circuit creation, simulating the circuits, and encoding each circuit into a tensor.

All of the 100-qubit simulations, data pre-processing, and model training were performed using two 14-core Intel Xeon CPU E5-2697 v3 @ 2.60GHz processors. An end-to-end run (i.e., circuit generation to trained model predictions) took roughly 12 hours of wall clock time.

B Code and data availability

The simulated data as well as records of all the quantum physics-aware networks will be released publicly once they clear Sandia’s copyright process. Until then, please email the authors. The CNNs and 5-qubit experimental datasets used in Hothem et al. [2024b] are available at Hothem et al. [2023a]. The datasets were originally located at Proctor et al. [2021b]. Each dataset was released under a CC-BY 4.0 International license.

All simulations were performed using a combination of `pygsti` version 0.9.11.2 [Nielsen et al., 2020] and `stim` version 1.13.0 [Gidney, 2021]. Models were trained and developed using Keras version 2.12.0 [Chollet et al., 2015] and TensorFlow version 2.12.0 [Abadi et al., 2015]. The physics-aware network model classes (`CircuitErrorVecScreenZErrorsWithMeasurementsBitstrings` for PST and `CircuitErrorVec` for process fidelity) are available in the Supplementary Material [Hothem et al., 2024a] as well as on the `feature – ml` branch of `pygsti`.

C Datasets

Device	Geometry	Circuit types	Circuit widths	Circuit depths	Training set size	Validation set size	Test set size
ibmq_london	t-bar	mirror	1-5 qubits	3-515 layers	711 circuits	104 circuits	91 circuits
ibmq_ourense	t-bar	mirror	1-5 qubits	3-515 layers	930 circuits	124 circuits	114 circuits
ibmq_essex	t-bar	mirror	1-5 qubits	3-515 layers	713 circuits	93 circuits	86 circuits
ibmq_burlington	t-bar	mirror	1-5 qubits	3-515 layers	682 circuits	90 circuits	92 circuits
ibmq_vigo	t-bar	mirror	1-5 qubits	3-515 layers	1029 circuits	137 circuits	126 circuits
ibmq_yorktown	bowtie	mirror	1-5 qubits	3-515 layers	1097 circuits	132 circuits	140 circuits
Ring (x5)	ring	random <i>i.i.d.</i>	1-4 qubits	1-180 layers	2813 circuits	938 circuits	1250 circuits
Ring (x5)	ring	mirror	1-4 qubits	8-174 layers	-	-	750 circuits
100-qubit Ring	ring	random <i>i.i.d.</i>	100 qubits	1-22 layers	2812 circuits	938 circuits	1250 circuits

Table 1: **Summary data of every dataset used in the paper.** The data for the 4-qubit ring processors is averaged over the 5 simulated datasets. See Figure 4 for images of each processor geometry (i.e., the qubit connectivity graph).

We provide additional details on the datasets used in the paper. Table 1 summarizes each dataset. We tracked all weight-2 errors with support on qubits connected by 2 hops in all, but the 100-qubit datasets. Below, we provide additional details on the circuit and error model generating processes.

C.1 Creating the circuits

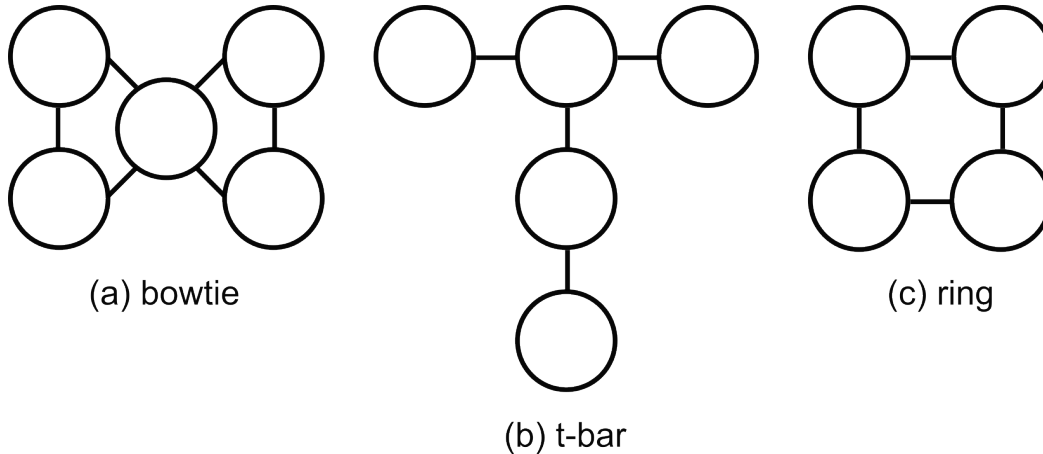


Figure 4: **Device geometries.** The connectivity graphs for the (a) 5-qubit ibmq_yorktown “bowtie” processor; (b) the remaining 5-qubit experimental “t-bar” processors; and (c) the 4-qubit simulated “ring” processor. The 100-qubit simulated “ring” processor has the same topology, just more qubits.

In this subsection, we go over how the random *i.i.d.*-layer circuits and random mirror circuits were created for this paper. We start by explaining how we generated the random *i.i.d.*-layer circuits for a 4-qubit ring processor, and then explain the modifications needed to generate the random mirror circuits. This subsection’s content is conceptual. The actual circuits were created in `pygsti` using the code in the Supplementary Material.

Each random *i.i.d.*-layer circuit c was created by a multi-step process. First, we randomly sampled a connected subset $\mathbb{Q}_c \subseteq \{Q0, Q1, Q2, Q3\}$ of qubits for which c is designed for. Then, we uniformly sampled c ’s depth from between 1 and d_w , a pre-determined, circuit-width-dependent maximum depth. The depths d_w were selected to ensure that $F(c) > 85\%$ given the maximum error strengths used to create the simulated error model (Section C.2). Third, we randomly sampled a two-qubit gate density ρ_{2Q} between 0 and $2/3$. The density ρ_{2Q} determines the average number of two-qubit gates in each of c ’s layers. We then sampled each layer *i.i.d.* from all possible circuit layers on the qubits in $\mathbb{Q}_c$.

The random mirror circuits were generated using a similar multi-step process with two differences. The first difference is that we used a pre-determined maximum depth of $d_w/6$. We chose to reduce the pre-determined, circuit-width-dependent maximum depth so that the deepest random mirror circuits had roughly the same length as the deepest random *i.i.d.* circuits. The second difference is that we created a random mirror circuit on $\mathbb{Q}_c$. See Proctor et al. [2021a] for more details.

C.2 Creating a 4-qubit error model

In this subsection, we explain how we constructed the 4-qubit Markovian local coherent error model used in Section 6. Again, we provide a conceptual explanation. The actual error model was created in `pygsti` using the code found in the Supplementary Material.

The 4-qubit Markovian local coherent error model was specified using the error generator framework explained in Section 2 and Blume-Kohout et al. [2022]. The error model consists of operation-dependent errors sampled according to a two-step process. The error strengths for each gate and qubit(s) pairs were independently sampled. First, we sampled an overall error strength ε_g for each one- and two-qubit gate g by randomly sampling from $[0, 1]$ and scaling by a pre-determined maximum error strength (.025%). Then we sampled the relative error strengths $\vec{\varepsilon}_{g,\text{rel}}$ of each of the $4^n - 1$ coherent errors, where $n = 1, 2$ for one- and two-qubit gates, respectively. We then normalized $\vec{\varepsilon}_{g,\text{rel}}$

to obtain the actual error strengths according to the following equation:

$$\vec{\varepsilon}_g = \frac{\sqrt{\varepsilon_g} \cdot \vec{\varepsilon}_{g,\text{rel}}}{\sqrt{\sum_i \varepsilon_{g,i}^2}}. \quad (8)$$

The re-scaling ensures that, to first order, gate g contributes approximately ε_g to the circuit's process infidelity (or PST, if appropriate).

C.3 Creating a 100-qubit error model

In this subsection, we explain how we constructed the 100-qubit Markovian error model used in Section 7. As with the 4-qubit error model, we provide a conceptual explanation. The actual error model was created in `pygsti`.

As with the 4-qubit error model, the 100-qubit Markovian error model was specified using the error generator framework explained in Section 2. The error model consisted of gate-dependent errors sampled according to a two-step process. Unlike with the 4-qubit error model, the 100-qubit error model included non-local Pauli stochastic and coherent errors, and the errors for each gate were independent of the qubit acted upon by the gate. Moreover, all of the errors in the 100-qubit error model are weight-1 errors (i.e., they affect a single qubit).

Each gate's error strengths were independently sampled. First, we enumerated all 600 possible weight-1 Pauli stochastic and coherent errors in a 100-qubit device. Then, for each gate g , we independently sampled the strength of each of the 300 weight-1 Pauli stochastic errors and the 300 possible weight-1 coherent errors. The strengths were sampled uniformly random, with a maximum Pauli stochastic error strength of 0.0000001 and a maximum coherent error strength of 0.00005. The resulting 600 error strengths were assembled into a 600-dimensional error vector for the gate, $\vec{\varepsilon}_g$.

C.4 Simulating the 100-qubit circuits

In this subsection, we describe the first-order simulation method used to approximate $F(c)$ in our 100-qubit simulations (Section 4.3). The method works by constructing an approximate error matrix $E(c)$ [Fig. 1(e)], and then estimating $F(c)$ by performing the same computation as in the second half of a qpa-NN [Fig. 1(f)]. Herein, we describe how to construct the approximate error matrix $E(c)$. Readers should refer to Section 3 for an in-depth explanation on how to use $E(c)$ to estimate $F(c)$.

For each circuit c , we constructed an approximate error matrix $E(c)$ by concatenating approximate error *vectors* $E_i(c)$ for each circuit layer in c . Each layer's error vector was computed as a linear combination of the individual gate error vectors, with coefficients equal to the number of times each gate appears in the circuit layer.

D Networks

D.1 Quantum-physics-aware network details

Dataset	Metric	Model size	N_{hops}	N_{errors}	Dense units
5-qubit t-bar	PST(c)	1218348	3	174	[30, 20, 10, 5, 5, 1]
5-qubit bowtie	PST(c)	1596420	3	210	[30, 20, 10, 5, 5, 1]
4-qubit ring	$F(c)$	299772	2	132	[30, 20, 10, 5, 5, 1]
100-qubit ring	$F(c)$	12706200	1	600	[30, 20, 10, 5, 5, 1]

Table 2: Summary data for the quantum-physics-aware networks used in the paper.

Table 2 briefly outlines the hyperparameters and model sizes of the physics-aware neural networks used in this paper. The N_{hops} and N_{errors} hyperparameters were chosen by hand based upon subject-matter-expert knowledge of the errors in a quantum computer. The size and shape of the dense layers

were selected arbitrarily. All dense subunits used a ReLU activation function. All models were trained using keras’s Adam [Kingma and Ba, 2015] optimizer with a step size of 10^{-3} and with mean squared error as the loss function. Model training was cut short using early stopping. To help with training, we scaled $PST(c)$ and $F(c)$ by 10000 when training the physics-aware networks. The notebooks in the Supplementary Material contain more details.

D.2 Convolutional neural network details

Details on the specific convolutional neural networks used in this paper are located in Hothem et al. [2024b]. We fine-tuned each network on high-PST experimental data using the Adam optimizer and early stopping.

We selected the model architecture for the convolutional neural networks used in our simulations via hyperparameter tuning. We performed 100 trials of hyperparameter tuning using the `BayesianOptimization` class in `kerastuner`. Additional details, including the specific hyperparameter space, are located in the Supplementary Material.

E Experimental results

Dataset	Network	Mean absolute error (%)	Bayes factor vs. CNN	Bayes factor vs. ft-CNN
ibmq_london	qpa-NN	1.09	10^{340}	10^{184}
	CNN	3.44	-	-
	ft-CNN	2.39	10^{156}	-
ibmq_ourense	qpa-NN	1.24	10^{30}	10^{33}
	CNN	1.55	-	-
	ft-CNN	1.68	$10^{-2.48}$	-
ibmq_essex	qpa-NN	1.39	10^{248}	10^{238}
	CNN	3.03	-	-
	ft-CNN	2.82	$10^{9.46}$	-
ibmq_burlington	qpa-NN	1.25	10^{152}	$10^{58.2}$
	CNN	2.27	-	-
	ft-CNN	1.61	$10^{93.7}$	-
ibmq_vigo	qpa-NN	1.21	10^{378}	10^{115}
	CNN	2.98	-	-
	ft-CNN	1.75	10^{263}	-
ibmq_yorktown	qpa-NN	1.19	10^{383}	$10^{28.4}$
	CNN	2.71	-	-
	ft-CNN	1.31	10^{354}	-

Table 3: **Summary model performance data on experimental data.**

Table 3 summarizes model performance on each of the experimental datasets used in the paper. Copies of the pre-trained quantum-physics-aware neural networks, original CNNs, and fine-tuned CNNs are available in the Supplementary Material. Scatter plots of each model’s predictions are also available in the Supplementary Material.

F Simulation results

Table 4 summarizes model performance on each of the 4-qubit simulated datasets used in the paper. Copies of the pre-trained quantum-physics-aware neural networks and fine-tuned CNNs are available in the Supplementary Material. Scatter plots of each model’s prediction errors are also available in the Supplementary Material.

Dataset	Network	Circuit type	Mean absolute error (%)	Pearson correlation coefficient
0	qpa-NN	random <i>i.i.d.</i>	.176	.968
		mirror	.720	.914
	CNN	random <i>i.i.d.</i>	.404	.751
		mirror	1.17	.872
1	qpa-NN	random <i>i.i.d.</i>	.200	.960
		mirror	.652	.921
	CNN	random <i>i.i.d.</i>	.421	.741
		mirror	.922	.861
2	qpa-NN	random <i>i.i.d.</i>	.190	.970
		mirror	.769	.914
	CNN	random <i>i.i.d.</i>	.406	.732
		mirror	1.20	.853
3	qpa-NN	random <i>i.i.d.</i>	.191	.952
		mirror	.719	.912
	CNN	random <i>i.i.d.</i>	.367	.764
		mirror	1.02	.857
4	qpa-NN	random <i>i.i.d.</i>	.195	.960
		mirror	.761	.898
	CNN	random <i>i.i.d.</i>	.405	.763
		mirror	1.05	.847

Table 4: **Summary model performance data on the 4-qubit simulated data.**

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Yes, we list the main claims and the section in which they are answered below.

- We developed a new physics-aware neural network architecture for quantum capability learning: Section 3.
- Our new approach achieves up to a $\sim 50\%$ and $\sim 76\%$ improvement over state-of-the-art convolutional neural networks on experimental and simulated data, respectively: Section 5 and 6.
- Our new approach beats state-of-the-art convolutional networks, in part, due to their improved ability to model coherent errors: Section 6.
- Our new approach achieves moderate prediction accuracy on an out-of-distribution prediction task: Section 6.
- Our approach scales to 100+ qubits.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: We include a discussion of the limitations of the work in Section 8.1.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: We do not include any new theoretical results or proofs.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: We believe that we provide sufficient details to reproduce the main experimental results of the paper. Readers should be able to recreate our results based on the details in the main body of the paper and the appendix, or by using the notebooks in the supplemental material once released.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: We plan to provide open access to all the data and code used in the paper through a GitHub repository. We also provide explicit instructions on how to access the experimental data in Appendix B.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.

- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: In Section 4 we explain how we processed the experimental and simulated datasets. Appendices C and D contain additional details on the datasets, specific network instantiations, and model training.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We report error bars on the mean absolute error of the models' predictions in Section 6. We do not report error bars on the models' predictions in the experimental data as the original paper [Hothem et al., 2024b] reported their error bars as being trivial. However, we do provide the standard deviation of the percent change in the MAE across the experimental datasets, and we report Bayes factors for each model on each experimental dataset, demonstrating substantial improvement by the physics-aware networks.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We provide details on the compute resources and compute time used in this work in Appendix A.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We believe that we have conducted our research in a manner that conforms, in every respect, with the NeurIPS Code of Ethics. The only relevant areas of concern are the use of deprecated datasets and respect for copyright and fair use. We believe that we have not violated either of these requirements, although the cloud-accessed processors used in the experimental section are no longer available (but the datasets are not deprecated).

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: We believe that there is little to no societal impact of our work. We believe that the following quote from the checklist guidelines is relevant: "it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster." While quantum computers might some day have a large societal impact, our work does not directly improve their ability to run programs with societal impact.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.

- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: Our paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We explicitly reference each existing dataset in the References, and state each existing datasets license in Appendix B. All existing assets were licensed under a CC-BY 4.0 license, requiring proper attribution.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: At the time of submission, we have released code used in Section 6 and the new simulated dataset in the supplemental material. In the future we will release additional assets, such as the models trained on experimental data, publicly after receiving approval from our employer.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.