

This is the accepted manuscript made available via CHORUS. The article has been published as:

Unsupervised quantum circuit learning in high energy physics

Andrea Delgado and Kathleen E. Hamilton

Phys. Rev. D **106**, 096006 — Published 9 November 2022

DOI: [10.1103/PhysRevD.106.096006](https://doi.org/10.1103/PhysRevD.106.096006)

Unsupervised Quantum Circuit Learning in High Energy Physics*

Andrea Delgado^{1,†} and Kathleen E. Hamilton^{2,‡}

¹*Physics Division, Oak Ridge National Laboratory, Oak Ridge, TN 37830*

²*Computer Science and Engineering Division, Oak Ridge National Laboratory, Oak Ridge, TN 37830*

Unsupervised training of generative models is a machine learning task that has many applications in scientific computing. In this work we evaluate the efficacy of using quantum circuit-based generative models to generate synthetic data of high energy physics processes. We use non-adversarial, gradient-based training of quantum circuit Born machines to generate joint distributions over 2 and 3 variables.

I. INTRODUCTION

High-energy physics seeks to understand matter at the most fundamental level. Vast and complex accelerators have been built to elucidate the dynamical basis of the fundamental constituents. At these large-scale facilities, high-performance data storage and processing systems are needed to store, access, retrieve, distribute, and process experimental data. Experiments like the Compact Muon Solenoid (CMS) and the ATLAS at the Large Hadron Collider (LHC) are incredibly complex, involving thousands of detector elements that produce raw experimental data rates over a Tb/sec, resulting in the annual productions of datasets in the scale of hundreds of Terabytes to Petabytes. In addition, the manipulation of these complex datasets into summaries suitable for the extraction of physics parameters and model comparison is a time-consuming and challenging task.

A crucial element of any analysis workflow in particle physics involves simulating the physical processes and interactions at these facilities to develop new theories and models, explain experimental data, and characterize background. These simulations also allow for studying detector response and plan detector upgrades. The simulation of particle interactions in the detector volume is often computationally intensive, taking up a significant fraction of the computational resources available to physicists.

Recently, alternative methods for detector simulation and data analysis tasks have been explored, like machine learning (ML) applications and quantum information science (QIS). Although ML applications have already been

incorporated in HEP experimental and data analysis environments, QIS is still an evolving field and its applications and their suitability remain to be explored. QIS is a rapidly developing field focused on understanding the analysis, processing, and transmission of information using quantum mechanical principles and computational techniques. QIS may be able to address the conventional computing gap associated with HEP-related problems, specifically those computational tasks that challenge CPUs and GPUs, such as efficient and accurate event generators. Thus, the relevance of having an efficient simulation mechanism that can faithfully reproduce particle interactions after a high-energy collision has sparked the development of alternative methods. One particular example is the use of generative models such as generative adversarial networks (GANs) [1], which have been utilized in HEP as a tool for fast Monte Carlo (MC) simulations [2–4], and as a machine learning-enhanced method for event generation [5–8]. Some of these studies report up to five orders of magnitude decrease in computing time. A crucial feature of generative models is their ability to generate synthetic data by learning from actual samples without knowing the underlying physical laws of the original system. In some studies, generative models have been shown to overcome the statistical limitations of an input sample in subtraction of negative-weight events in samples generated beyond leading-order in QCD [5], and to increase the statistics of centrally produced Monte Carlo datasets [7]. Quantum-assisted models have also been proposed for Monte Carlo event generation [9], detector simulation [10, 11], and determining the parton distribution functions in a proton [12]. In this work, we successfully trained a quantum generative model to reproduce kinematic distributions of particles in pp interactions at the LHC, with high fidelity. Quantum circuit Born machines (QCBM) trained via gradient-based optimization [13–15] are examples of circuit-based parameterized models that can be trained on near-term quantum platforms.

II. MODEL AND LEARNING ALGORITHM

Although generative models trained in adversarial settings have been proven to be a valuable tool in HEP, for this work, we focus on generative models trained

* This manuscript has been authored by UT-Battelle, LLC, under Contract No. DE-AC0500OR22725 with the U.S. Department of Energy. The United States Government retains and the publisher, by accepting the article for publication, acknowledges that the United States Government retains a non-exclusive, paid-up, irrevocable, world-wide license to publish or reproduce the published form of this manuscript, or allow others to do so, for the United States Government purposes. The Department of Energy will provide public access to these results of federally sponsored research in accordance with the DOE Public Access Plan.

[†] delgadoa@ornl.gov

[‡] hamiltonke@ornl.gov

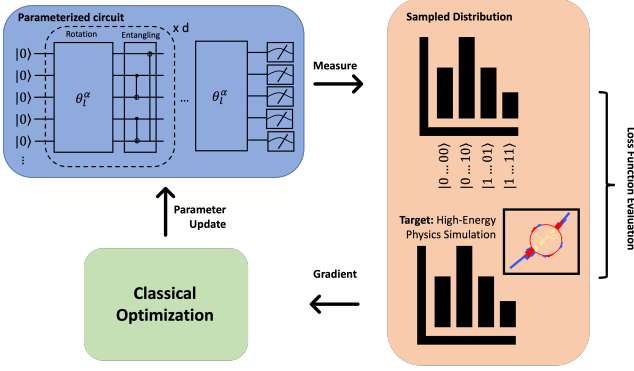


FIG. 1: Diagram of the differentiable QCBM training scheme.

with non-adversarial methods. Expressly, we set up and train multiple QCBM using data-driven circuit learning (DDCL) [13–15]. DDCL employs a classical-quantum hybrid feedback loop, as described in Fig. 1.

Generating synthetic data for this HEP application consists of three stages: first, we encode our data of M observations with N -dimensional, numerical features ($\mathcal{D} = \{X^{(1)}, \dots, X^{(m)}, \dots, X^{(M)}\}$) into a distribution $P(x)$ defined over finite length bitstrings (x); second, we train a parameterized circuit ansatz using DDCL to prepare an approximation of this distribution $\tilde{P}(x)$; third, we generate synthetic data by decoding bitstrings sampled from $\tilde{P}(x)$. In the following subsection, we describe these stages in greater detail.

A. Data Encoding

The kinematic distributions of a particle jet in a high-energy collider experiment such as the LHC can be constructed as a 2-dimensional joint distribution (over $(p_T, mass)$ features) or 3-dimensional joint distributions (over $(p_T, mass, \eta)$ features). Each of these kinematic variables is encoded as a discretized binary string using $q = Q/N$ qubits. The marginal distribution is discretized into 2^q bins and each bin index is converted to a 2^q -length binary bitstring.

For a QCBM constructed with Q total qubits, $P(x)$ is constructed by concatenating lists of binary bitstrings which encode the marginals of individual features in a classical dataset. The N -dimensional correlated data features ($X^{(m)}$) are encoded as 2^q -length binary strings (x_i). Thus, the total Q qubits contain a concatenation of q -dimensional bitstrings. The final 2^q -length bitstrings are constructed by concatenating the N feature 2^q -length binary bitstrings and normalizing the amplitudes.

B. Quantum Circuit Model and Training

The QCBM is an example of an implicit model for generative learning [16] that generates data by measuring the system as a Born machine. A QCBM is a parameterized unitary $\mathcal{U}(\Theta)$ that prepares Q -qubits in the state $|\psi_\Theta\rangle = \mathcal{U}(\Theta)|\psi_0\rangle$. The initial state $|\psi_0\rangle$ is fixed, and in this study, we use: the all zero-state $|\psi_0\rangle = |0\rangle^{\otimes Q}$; a product state of $Q/2$ Bell states $|\Phi^+\rangle^{\otimes Q/2}$; and a product state of $Q/3$ GHZ states $|\psi_0\rangle = |\text{GHZ}\rangle^{\otimes Q/3}$.

Measuring $|\psi_\Theta\rangle$ in a fixed basis $\mathcal{M}$ ¹ requires sampling from the state with N_{shots} shots. This defines a classical distribution over the 2^Q computational basis states $\tilde{P}_\Theta(x)$ that is used in training, and later used to generate the synthetic data $\tilde{X}$.

1. Parameterized Quantum Circuit

Finding the ideal unitary $\mathcal{U}(\Theta)$ is dependent on the parameterized quantum circuit (PQC). The PQC design plays an essential role in the performance of many variational hybrid quantum-classical algorithms [17, 18] by defining the hypothesis class. For our application, PQCs must be able to model different types of correlations in the input data. This requires circuits to prepare strongly entangled quantum states and the ability to explore Hilbert space.

Variational algorithms have been implemented using quantum circuits composed of a network of single and two-qubit operations, with rotation angles serving as variational parameters. The pattern defining the network of gates is referred to as a *unit-cell* or *circuit block* that can be repeated to suit the needs of the application or task at hand. Recently, the term Multilayer Quantum Circuit (MPQC) was coined to describe this type of variational circuit architecture [19]. In this study, we train two circuit templates or ansatz. Each circuit is defined using a Q -qubit register and specified by the number of layers d , consisting of a rotation and an entangling component. Ansatz 1 has a combined rotation and entangling gates in a “Brick Layer” or “Simplified 2-design” architecture and has been shown to exhibit important properties to study barren plateaus in quantum optimization landscapes [20, 21]. Ansatz 2 was chosen due to the low correlation displayed between variables and is an extension of ansatz employed in benchmarking tasks [14].

The diagram for one layer of each template is shown in Fig. 2. The rotation gate layers are parameterized by the arbitrary single-qubit rotation gate implemented in PennyLane [22], which has 3 rotation angles $R(\Theta)_\ell^{(i)} = R(\omega, \theta, \phi)_\ell^{(i)}$, where the layer index ℓ runs from 0 to d , and i is the qubit index.

¹ We use the Z-basis $\mathcal{M} = Z_i \otimes \dots \otimes Z_Q$ unless otherwise noted.

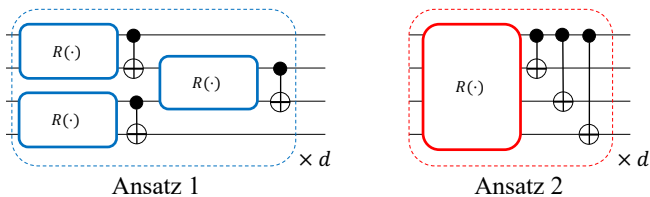


FIG. 2: Diagram for one layer of the circuit templates used to construct and train the QCBM. Both ansatz are constructed using a layered pattern of rotation and entangling gates. For both ansatz a final layer of rotation gates is added before measurement.

We choose to train two ansatz that can be embedded onto near-term NISQ devices: either using a 1D chain of qubits (for Ansatz 1), or the "disconnected tree" configuration (for Ansatz 2). For Ansatz 1 used in this study, the number of parameters to optimize during training is $N_{\text{par}}(N_\ell, Q) = 3N_\ell[2Q - 1(2)] + 3Q$ for QCBM with odd(even) qubits Q . For Ansatz 2, $N_{\text{par}}(N_\ell, Q) = 3(N_\ell + 1)Q$, regardless of whether Q is even or odd.

2. Training and cost function

Training a QCBM via DDCL optimizes Θ by minimizing a loss function $\mathcal{L}(P, \tilde{P}_\Theta)$ such that $\tilde{P}_\Theta(x) \approx P(x)$. We use the Jensen-Shannon (JS) divergence- a differentiable function that compares two distributions $P(x)$ and $\tilde{P}_\Theta(x)$:

$$JS(P|\tilde{P}_\Theta) = \frac{1}{2} \sum_x \left[P \log \left(\frac{P}{M} \right) + \tilde{P}_\Theta \log \left(\frac{\tilde{P}_\Theta}{M} \right) \right] \quad (1)$$

where $M = (P + \tilde{P}_\Theta)/2$. The minimum value of the loss function $JS(P|\tilde{P}_\Theta) = 0$ is achieved when $\tilde{P}_\Theta(x) = P(x)$. The model parameters Θ are optimized using classical gradient descent methods so that the loss function is minimized. The gradient of the loss function is computed with respect to the circuit parameters using the parameter shift rule [23]. Each QCBM is trained using the Adam gradient-based optimizer [24] available in the PennyLane library [22].

C. Measurement decoding and post-processing

The output of a QCBM is $\tilde{P}_\Theta(x)$: a classical distribution over 2^Q -length binary strings. To convert $\tilde{P}_\Theta(x)$ to numerical n-dimensional features ($\tilde{\mathcal{D}} = \{\tilde{X}^{(1)}, \dots, \tilde{X}^{(m)}, \dots\}$), we reverse the steps described in Section II A: each binary 2^Q length string is disassociated into N composite strings each of length 2^q , and a float value randomly drawn from a uniform distribution

defined with the bin edges previously used to map the samples x_m into the binary basis to generate $P(x)$.

III. VALIDATION ON LHC DATASET

One of the big computational challenges in HEP is the considerable computing time required to model the behavior of subatomic particles both at the vertex and detector level. In Section II, we introduce the architecture of a quantum generative model that aims to provide an alternative to traditional Monte Carlo (MC) methods in the context of data augmentation. Thus, to validate the proposed model, we consider the simulation of the production of pairs of jets in pp interactions at the LHC. The dataset [6] consists of 10 million di-jet events generated using MADGRAPH5 version 2.6.4 [25] and PYTHIA8 version 8.307 [26], corresponding to a center of mass energy of 13 TeV and an integrated luminosity of about 0.5 fb^{-1} . The response of the detector was simulated by a DELPHES version 3.4.3pre06 [27] fast simulation, using settings that resemble the ATLAS detector. An average of 25 additional soft-QCD pp collisions (pile-up) were added to the simulation to mimic the conditions of a typical collider event realistically. Jets were reconstructed using the anti- k_T [28] algorithm as implemented in FASTJET [29], with a distance parameter $R = 1.0$. A selection cuts on the scalar sum of the transverse momenta of the outgoing partons $HT \geq 500 \text{ GeV}$ was applied, reducing the sample size to about 4 million events. The kinematic distributions of the leading jet on the di-jet system are used to validate the QCBM models and evaluate its performance in a real-world application.

IV. RESULTS

In this section, we report on the results of training a QCBM to prepare the target distribution of the dataset described in Section III. The model performance is evaluated by studying the JS divergence value throughout the training and comparing target (MC expectation) and generated (the output of the trained QCBM model) distributions. We explore the encoding of the target distributions for 2(3) variable joint distributions into 8(12) qubit systems. This scheme allows for a four-qubit encoding per distribution. Unless noted, each marginal distribution is encoded in a target state binned over $2^4 = 16$ basis states. We trained the QCBM circuits in the absence of noise to obtain a set of optimal parameters Θ . Then, the circuits were deployed using the trained parameters on IBM quantum devices to study the effect of noise in the loss landscape, reproducing the target distribution. Finally, a local parameter tuning scheme was applied to improve the performance in the presence of noise.

A. Training with noiseless qubits

1. 2D Distributions

In Figure 3, the JS divergence values are plotted as a function of training step. The circuits were trained using $N_{shots} = 8192$ to prepare a joint 2D distribution corresponding to the marginal distributions of the leading jet transverse momentum (p_T) and mass, both binned over the 16 basis states corresponding to four qubits. Circuits were constructed using the ansatz configurations shown in Figure 2 and trained to start from either the all-zero state ($|\psi_0\rangle = |0\rangle^{\otimes 8}$), or from a product of four Bell states ($|\Phi^+\rangle^{\otimes 4}$). We fixed the number of layers $N_{layers} = 6$ for Ansatz 1 (blue) and 12 for Ansatz 2 (red). Each circuit is trained for 300 steps of Adam with a learning rate $\alpha = 0.01$.

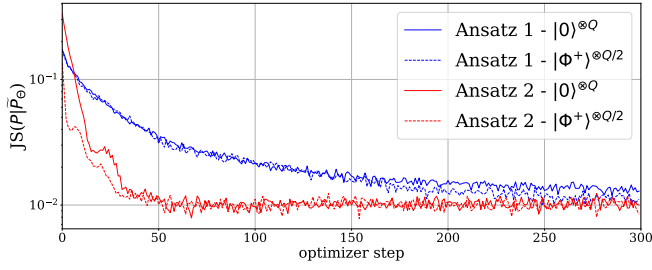


FIG. 3: JS divergence value as a function of training step using $N_{shots} = 8192$. Blue(red) points correspond to Ansatz 1(2). Circuits were initialized in the all zero state (solid lines) or four Bell states (dashed lines) and trained to learn a 2D joint distribution.

From Figure 3 we can conclude that: Ansatz 2 converges to a stable JS divergence value much faster than Ansatz 1, and the training of the QCBM is not affected by the choice of the initial state. Figure 4 displays the distributions of samples generated via projective measurements on the qubits in the trained circuits. We observe that the data generated resembles the target distributions with high fidelity, with a slightly better agreement for data generated by sampling from the QCBM constructed with Ansatz 2 (red).

In Figure 5, we compare the sampled and target distributions obtained when starting the training from either an all-zero state ($|\psi_0\rangle = |0\rangle^{\otimes 8}$) or from a product of four Bell states ($|\Phi^+\rangle^{\otimes 4}$). We define a similarity measure to perform a systematic comparison of the marginal distributions for each feature by computing the mean absolute error (MAE) per bin between all normalized target and sampled marginal distributions:

$$D(p||\tilde{p}(\Theta)) = \frac{1}{N2^q} \sum_n \sum_i^{2^q} |p_i^{(n)} - \tilde{p}_i^{(n)}(\Theta)| \quad (2)$$

From both Figure 5 and Table I, we can conclude that Ansatz 2, initialized in the all-zero state, reproduces the

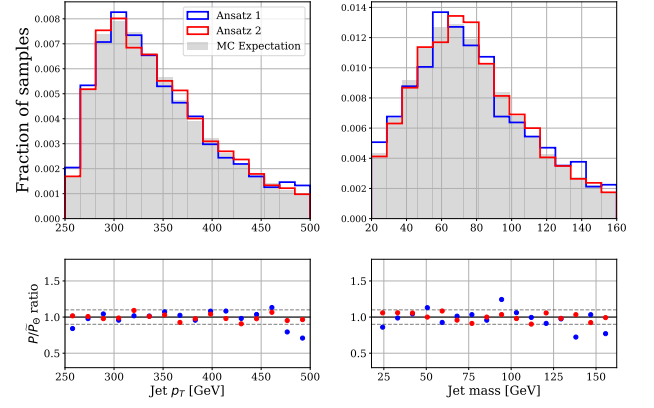


FIG. 4: **Top:** Target and sampled distributions. Blue(red) points correspond to Ansatz 1(2). Circuits were initialized in the all-zero state ($|\psi_0\rangle = |0\rangle^{\otimes 8}$) and trained to learn a 2D joint distribution. **Bottom:** Ratio of target and sampled distributions with horizontal guide lines marking $P/\tilde{P}_\Theta = 0.9$ and $P/\tilde{P}_\Theta = 1.1$.

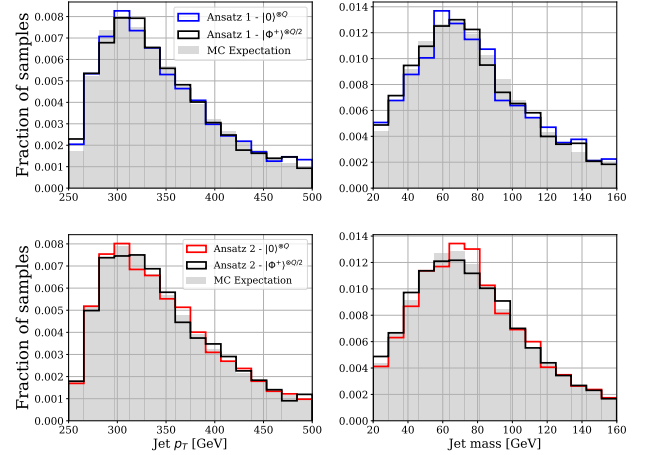


FIG. 5: **Top:** Target and sampled distributions constructed using Ansatz 1. Circuits were initialized in the all-zero state (blue) and four Bell states (black). **Bottom:** Target and sampled distributions constructed using Ansatz 2. Circuits were initialized in the all zero-state (red) and four Bell states (black).

feature marginals with the highest fidelity. On the other hand, the difference in $D(p||\tilde{p}(\Theta))$ for the two initial configurations considered is negligible, considering a statistical error proportional to $1/\sqrt{N_{samples}} \sim 0.0005728$, implying that the training is independent of circuit initialization.

Another important factor in the process of building the circuit to prepare the target state is the number of layers the template or ansatz in Figure 2 is repeated. This choice will also determine the number of trainable parameters in the model. In Figure 6, the JS divergence value is plotted as a function of training steps for $d = \{1, 3, 6, 9\}$ layers in Ansatz 1; and $d = \{2, 6, 12, 18\}$ layers in Ansatz

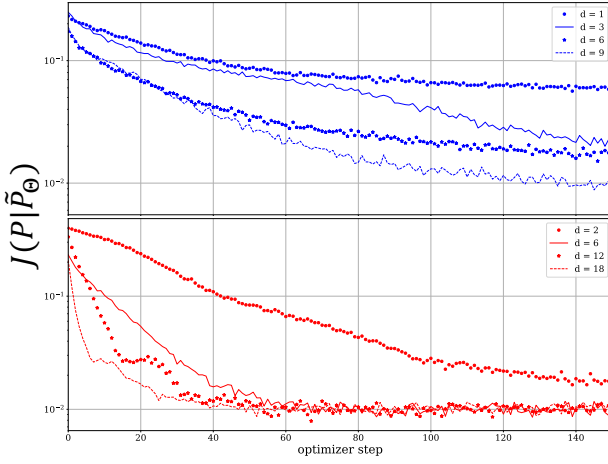


FIG. 6: Loss as a function of training step using $N_{shots} = 8192$. The different curves represent a different number of layers d , used to construct the 8 qubit QCBM. **Top:** Ansatz 1 (blue), and **Bottom:** Ansatz 2 (red).

Ansatz	N_{par}	Initial State	$D(p \tilde{p}(\Theta))$
1	276	$ \psi_0\rangle = 0\rangle^{\otimes 8}$	0.007153
1		$ \psi_0\rangle = \Phi^+\rangle^{\otimes 4}$	0.006767
2	312	$ \psi_0\rangle = 0\rangle^{\otimes 8}$	0.005452
2		$ \psi_0\rangle = \Phi^+\rangle^{\otimes 4}$	0.005696

TABLE I: $D(p|\tilde{p}(\Theta))$ for 8 qubit QCBM.

2. We can see from this plot that the number of layers used to construct Ansatz 2 has little impact on the minimal JS value reached after the training converged. Nonetheless, it affects how fast the model reaches this minimal JS value. On the other hand, the training performance of Ansatz 1 is highly dependent on the number of layers used to construct the circuit. We chose to use $d = 6(12)$ to construct our QCBM with Ansatz 1(2) to keep an optimal balance between training time and performance.

This study proposes using quantum generative models as a data augmentation tool. The ability of generative models in increasing the statistical precision of the simulated events beyond the training sample has been studied in Ref. [30, 31]. The results presented in this section thus far were obtained by training QCBMs on a target distribution using close to 4 million events. We also investigated how reducing the training dataset size affects the trained model's fidelity to reproduce the target distribution. In Figure 7, the horizontal axis represents the fraction of the initial training dataset of 4 million events used to train the QCBM. Once the model is trained, using a fraction of the full dataset, the JS divergence metric is evaluated on the target distribution generated using the whole dataset. The distribution was obtained by evalu-

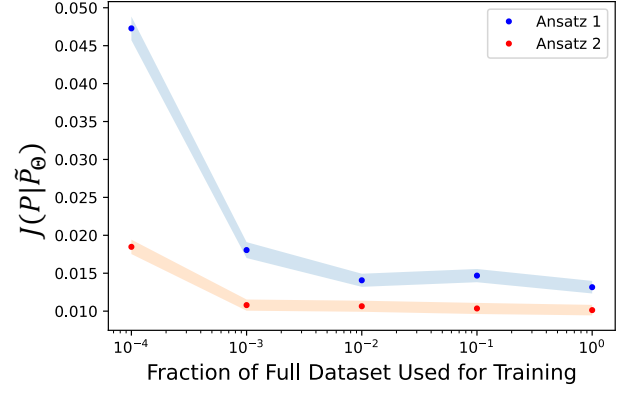


FIG. 7: Mean loss between the complete dataset and the distribution generated by 8-qubit QCBM models QCBM circuits plotted as solid circles in blue(red) for Ansatz 1(2). Each QCBM model was trained on partial data. Shaded regions correspond to mean JS divergence value $\pm$ one standard deviation.

ating the QCBM with the trained parameters with 8192 shots. The sampling process was repeated 1000 times, and the mean is reported in Figure 7 as a solid blue(red) dot for Ansatz 1(2). The bands correspond the mean JS divergence value $\pm \sigma$.

(a) Monte Carlo (Ground Truth)

	p_T	mass
p_T	-	0.2
mass	0.2	-

(b) Ansatz 1

	p_T		mass	
	$ 0\rangle^{\otimes 8}$	$ \Phi^+\rangle^{\otimes 4}$	$ 0\rangle^{\otimes 8}$	$ \Phi^+\rangle^{\otimes 4}$
p_T	-		0.19	<i>0.12</i>
mass	0.19	<i>0.12</i>	-	

(c) Ansatz 2

	p_T		mass	
	$ 0\rangle^{\otimes 8}$	$ \Phi^+\rangle^{\otimes 4}$	$ 0\rangle^{\otimes 8}$	$ \Phi^+\rangle^{\otimes 4}$
p_T	-		-1.0e-3	<i>-9.1e-3</i>
mass	-1.0e-3	<i>-9.1e-3</i>	-	

TABLE II: Correlation matrices between jet p_T and mass (m) variables in the (a) target distribution, and samples obtained from the evaluation of the QCBMs constructed using (b) Ansatz 1 and (c) Ansatz 2 in Figure 2 with the trained parameters. Values displayed for initial states prepared in the all-zero state (bold) and a product of 4 Bell states (italics).

Finally, we report on the correlation matrix between the jet p_T and mass variables used to construct the target distribution. In IIa, the values associated with the target distribution are displayed in black. If the trained QCBM learned the joint distribution, one would expect to re-

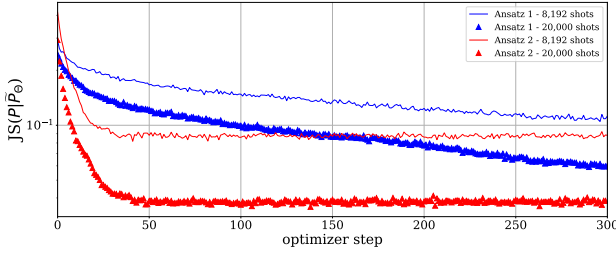


FIG. 8: JS divergence as a function of training step using $N_{shots} = 8192$ (solid lines) and $N_{shots} = 20,000$ (triangles). Blue(red) points correspond to Ansatz 1(2). Circuits were initialized in the all zero state and trained to learn a 3D joint distribution. QCBM were initialized in the all-zero state ($|\psi_0\rangle = |0\rangle^{\otimes 12}$).

cover the correlation matrix when evaluating the QCBM with the trained parameters. The results in Table II indicate that this is true for the QCBM constructed with Ansatz 1 (red), but not for the QCBM constructed with Ansatz 2. The correlation matrix for the latter case indicates that there is little correlation between the marginal distributions in the synthetic samples.

2. 3D Distributions

To understand how the trainability of non-adversarial generative models scales with the number of quantum registers, we increased the number of qubits in our model from 8 to 12. This increment translates into a larger number of basis states ($2^8 = 256$ to $2^{12} = 4096$). Furthermore, the joint probability distribution that we encode in the target state is now three-dimensional by including an additional marginal distribution associated with the "forwardness" of the jet with respect to the beam (jet η). In Figure 8, the JS divergence loss is plotted as a function of training step. The circuits were initialized in the all-zero state ($|\psi_0\rangle = |0\rangle^{\otimes 12}$) and $d = 6(12)$ to construct our QCBM with Ansatz 1(2). In this plot, we can also see the effect of increasing the number of shots during training, reporting a significant difference in JS divergence values when increasing N_{shots} from 8,192 to 20,000.

k When training QCBM with $Q = 12$ qubits, we also used an initial state ($|\psi_0\rangle = |\text{GHZ}\rangle^{\otimes Q/3}$) where 3-qubit subsets are initialized in a GHZ state. The comparison for the JS divergence value as a function of the training step from the three initial states considered is displayed in Figure 9. For Ansatz 1 (blue), the JS value for the three configurations is very similar for the first 100 training steps. From step 100 on, the QCBM initialized in a product of Bell states (blue cross) trains faster and converges to a lower JS value than the QCBMs initialized in the all-zero (solid blue) and all-plus (blue star) state. For Ansatz 2, the effect of the initial state used in the preparation of the circuit has a more negligible effect on

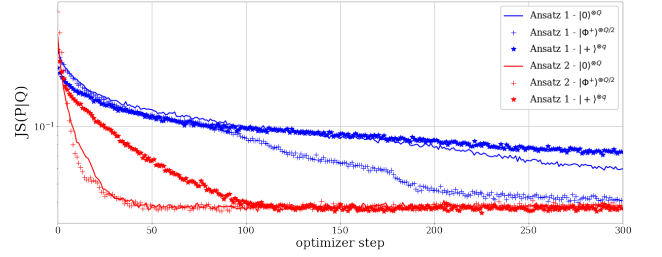


FIG. 9: JS divergence as a function of training step using $N_{shots} = 20,000$. Blue(red) points correspond to Ansatz 1(2). Circuits were initialized in the all-zero state (solid lines), a product of Bell states (cross), or the all-plus state (stars), and trained to learn a 3D joint distribution.

the minimal JS value the model converges after training. Nonetheless, the QCBMs prepared from the all-plus state (stars) seem to take longer to converge.

Ansatz	N_{par}	Initial State	$D(p \tilde{p}(\Theta))$
1	432	$ \psi_0\rangle = 0\rangle^{\otimes 12}$	0.0226
1		$ \Phi^+\rangle^{\otimes 6}$	0.0138
1		$ \psi_0\rangle = \text{GHZ}\rangle^{\otimes 3}$	0.0187
2	468	$ \psi_0\rangle = 0\rangle^{\otimes 12}$	0.0108
2		$ \Phi^+\rangle^{\otimes 6}$	0.0090
2		$ \psi_0\rangle = \text{GHZ}\rangle^{\otimes 3}$	0.0106

TABLE III: Comparison between target and sampled distributions according to Eq. 2.

In Figure 10, we compare the distributions of samples generated via projective measurements on the circuits evaluated on the trained parameters. The circuits were initialized in the all-zero state. By looking at Figure 10 and Table III, we observe a degraded performance in terms of similarity metric (Eq. 2) when compared to the 8-qubit circuit results. The $D(p|\tilde{p}(\Theta))$ value increases from 0.007153 to 0.02226 for Ansatz 1 and from 0.005452 to 0.0108 for Ansatz 2. Again, the trained QCBM that prepares the target distribution with the highest fidelity is Ansatz 2.

In Figure 11, we compare the distributions generated by QCBMs initialized in the three different initial configurations. Again, we see little dependence on the initial state for QCBMs prepared using Ansatz 2. Nonetheless, the effect of the initial state in QCBMs prepared using Ansatz 1 is now more evident.

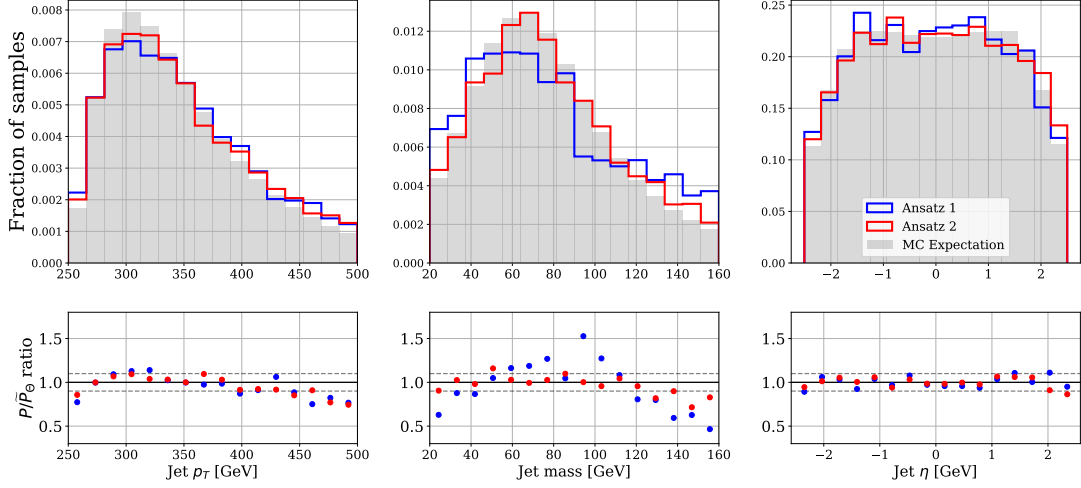


FIG. 10: **Top:** Target and sampled distributions. Blue(red) points correspond to Ansatz 1(2). Circuits were initialized in the all zero state ($|\psi_0\rangle = |0\rangle^{\otimes 12}$) and trained to learn a 3D joint distribution. **Bottom:** Ratio of target and sampled distributions with horizontal guide lines marking $Q/P = 0.9$ and $Q/P = 1.1$.

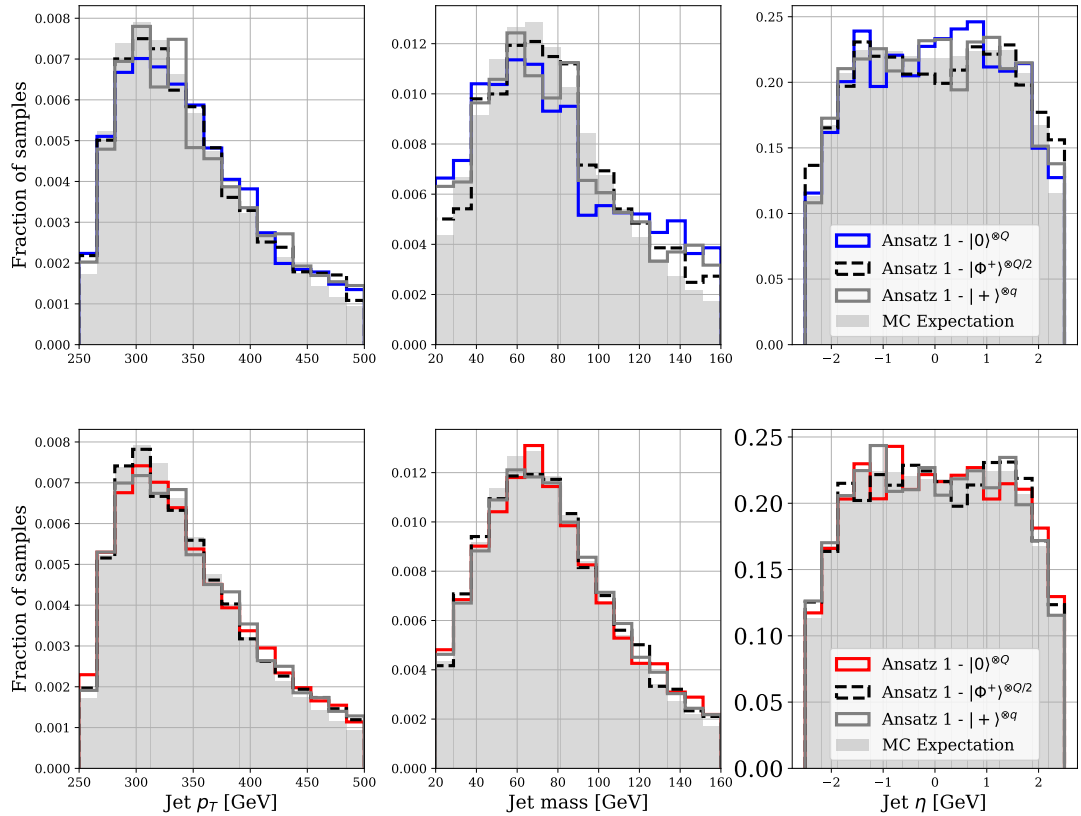


FIG. 11: **Top:** Target and sampled distributions corresponding to Ansatz 1. Circuits were initialized in the all zero state (solid blue line), a product of Bell states (black dashed line), and the all-plus state (gray dashed line) and trained to learn a 3D joint distribution. **Bottom:** Target and sampled distributions corresponding to Ansatz 2. Circuits were initialized in the all zero state (solid red line), a product of Bell states (black dashed line), and the all-plus state (gray dashed line).

(a) Monte Carlo (Ground Truth)

	p_T	mass	η
p_T	-	0.2	$7.3e-12$
mass	0.2	-	$2.7e-11$
η	$7.3e-12$	$2.7e-11$	-

(b) Ansatz 1

	p_T			mass (m)			η		
	$ 0\rangle^{\otimes 12}$	$ \Phi^+\rangle^{\otimes 6}$	$ \text{GHZ}\rangle^{\otimes 3}$	$ 0\rangle^{\otimes 12}$	$ \Phi^+\rangle^{\otimes 6}$	$ \text{GHZ}\rangle^{\otimes 3}$	$ 0\rangle^{\otimes 12}$	$ \Phi^+\rangle^{\otimes 6}$	$ \text{GHZ}\rangle^{\otimes 3}$
p_T	-	-	-	0.16	<i>0.16</i>	0.18	8.2e-3	<i>1.2e-3</i>	<i>-1.3e-3</i>
m	0.16	<i>0.16</i>	0.18	-	-	-	-0.014	<i>4.4e-3</i>	<i>7.7e-3</i>
η	8.2e-3	<i>1.2e-3</i>	<i>-1.3e-3</i>	-0.014	<i>4.4e-3</i>	<i>7.7e-3</i>	-	-	-

(c) Ansatz 2

	p_T			mass (m)			η		
	$ 0\rangle^{\otimes 12}$	$ \Phi^+\rangle^{\otimes 6}$	$ \text{GHZ}\rangle^{\otimes 3}$	$ 0\rangle^{\otimes 12}$	$ \Phi^+\rangle^{\otimes 6}$	$ \text{GHZ}\rangle^{\otimes 3}$	$ 0\rangle^{\otimes 12}$	$ \Phi^+\rangle^{\otimes 6}$	$ \text{GHZ}\rangle^{\otimes 3}$
p_T	-	-	-	-4.1e-3	<i>4.6e-3</i>	0.019	-3.7e-3	<i>-9.1e-3</i>	<i>1.6e-3</i>
m	-4.1e-3	<i>4.6e-3</i>	0.019	-	-	-	6.3e-3	<i>4.9e-3</i>	<i>2.2e-3</i>
η	-3.7e-3	<i>-9.1e-3</i>	<i>1.6e-3</i>	6.3e-3	<i>4.9e-3</i>	<i>2.2e-3</i>	-	-	-

TABLE IV: Correlation matrices between jet p_T , mass (m), and η variables in the (a) target distribution, and samples obtained from the evaluation of the QCBMs constructed using (b) Ansatz 1 and (c) Ansatz 2 in Figure 2 with the trained parameters. Values displayed for initial states prepared in the all-zero state (bold), a product of 6 Bell states (italics), and a product of 3 GHZ states.

Finally, we report on the correlation matrix between the jet p_T , mass and η variables used to construct the target distribution. In Table IV, the values associated with the target distribution are displayed in black. The results in Table IV display a slight discrepancy in the original correlation matrix (training dataset) and that for the samples generated by sampling Ansatz 1 (blue) with the trained parameters. Again, for the QCBM constructed with Ansatz 2, the matrix values indicate that there is little correlation between the marginal distributions in the synthetic samples.

B. Noisy Training with Layer-wise Coordinate Descent

The gradient-based training of the QCBM models presented in IV A was executed on noiseless (ideal) qubits but was not used for training on noisy hardware. However, the performance of the trained QCBM model on near-term quantum devices will be heavily impacted by hardware noise. Qubit initialization, gate noise, and errors in the circuit measurement step, all result in state preparation error that can cause parameterized models to converge to maximally mixed states, with an overall effect of flattening the loss landscape [32] and manifests in a noisy estimation of $\tilde{P}_\Theta(x)$.

While developing error mitigation methods that can be incorporated into variational training algorithms is an open area of research, one approach to noise mitigation is to implement the circuit training with hardware noise in order to learn optimized parameters that can compensate for time-independent errors, such as over- and under-rotation in single qubit gates. We test the robustness of the final parameters found in Section IV A

to hardware noise by first executing the trained QCBM models constructed with Ansatz 2 on superconducting qubit devices. The 8-qubit QCBM (276/312 parameters) was executed on the 16-qubit device `ibmq_guadalupe` and the 12-qubit QCBM (468) was executed on the 27-qubit device `ibmq_cairo`, both were accessed through a cloud-based queue.

In the absence of error and noise mitigation we instead opted to use a localized search over individual parameters. Each parameterized Ansatz (shown in Fig. 2) are constructed with layers of parameterized rotation gates, implemented using an arbitrary unitary gate with 3 rotational parameters. Then, we used a layer-wise coordinate descent (LCD) method to optimize each QCBM performance on hardware. The workflow is shown in Fig. 12. No readout error mitigation or other noise mitigation methods were used. LCD optimizes the parameters of a circuit $\mathcal{U}(\Theta)$ by searching the multi-dimensional parameter space along linear cuts of finite width. With N qubits in the register, each parameter was swept through a shift of parameters defined by $\epsilon = -\pi/4$ and spacing $2\epsilon/n$. The targeted backends allowed for a maximum number of circuits per batch (B) which defines the spacing $n = B/N$. For the 8-qubit QCBM trained on `ibmq_guadalupe` this resulted in a grid spacing of 0.0419. For the 12-qubit QCBM trained on `ibmq_cairo` this resulted in a mesh spacing of 0.0628. Each circuit was sampled using $N_{shots} = 20000$. The rotational parameters are optimized starting with the gates closes to the measurement process.

The top of Figure 13 displays the quartiles of JS loss over each iteration of LCD with 8-qubit QCBM circuits. The quartiles are plotted for the QCBM constructed with Ansatz 1 and 2 and evaluated with the updated parameters after each iteration on the `ibmq_guadalupe` backend in blue and red, respectively. The plot also shows how the JS divergence value degrades as the LCD training parameters deviate from those obtained during the noiseless training. On the other hand, hardware performance is relatively stable. The bottom of Figure 13 shows the sampled distributions generated by the evaluation of the QCBM with the parameters that yielded the lowest JS divergence value during the LCD training. We observe a more significant discrepancy between the target and sampled distributions compared to the results reported in Figure 4, where the QCBM is evaluated with the parameters obtained during the noiseless training. In Figure 14, the quartiles of JS loss (top) and the sampled and target distributions (bottom) are displayed. The QCBMs are constructed using Ansatz 2, to prepare a 3D joint distribution in a 12-qubit register. The LCD training was performed on the `ibmq_cairo` device.

In Figure 15 we show the quartiles of JS for three iterations representative of the different stages of the LCD training (left). During the first iterations in the training, the loss landscape is very flat, and the JS value is lower for the QCBMs evaluated on the `qasm_simulator`. The box plots on the left correspond to the JS values obtained

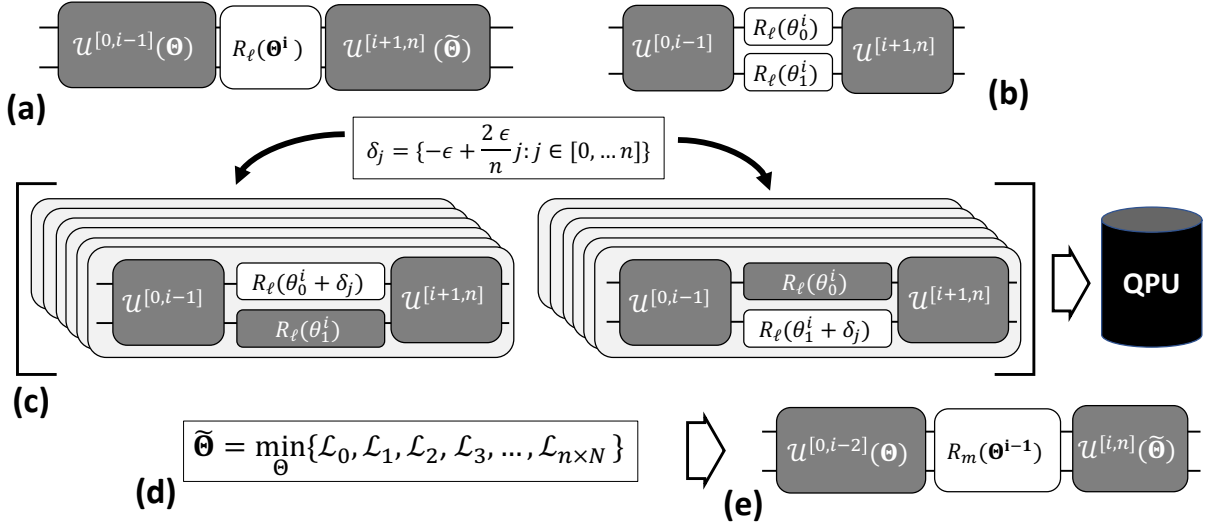


FIG. 12: Layer-wise coordinate descent (LCD) workflow. (a) a single rotational layer in a N -qubit QCBM model, (b) is composed of rotational gates R_ℓ acting on individual qubits. A batch of $n \times N$ circuits is constructed by sweeping each individual gate over a discrete set of n shifts and executed on a quantum processor (QPU). The loss is evaluated for each circuit executed (d), the parameter vector Θ is updated by the values which return the minimal value, and the updated parameter vector is used to start the search over the next rotational layer (e).

during the parameter sweep in the range $\epsilon \in [-\frac{\pi}{4}, \frac{\pi}{4}]$ for a particular rotation gate in a given layer. Then, we observe that, for subsequent iterations, the parameter tuning moves the parameters into regions where the performance of the noiseless simulator is degraded, and slowly improving the performance on the quantum device.

As we can see from Figures 13 and 14, noisy training can improve performance, but there are a number of obstacles. When the training is initialized with parameters pre-trained with noiseless qubits, multiple interactions may be needed to re-train. The gradual improvements in the loss function may not be robust against large deviations in the device noise. For example, executing the LCD workflow over all rotational parameters of the 12-qubit QCBM was done over multiple days, probably causing the discrete change in the loss between iterations 31 and 32.

V. QCBM DESIGN SPACE

Parameterized circuit ansatz used for variational algorithms need to balance expressability with trainability. Characterizing expressability is a difficult problem, and most common analysis relies on computing frame potentials which compare the distributions of states that a particular ansatz can generate, to the Haar random distribution [17, 18]. There are considerable ongoing efforts in determining the characteristics of circuits that lead to effective training and scaling [33]. In this paper we used two different ansatz designs and multiple initializations for the Q -qubit register. In this section we discuss some

observations based on our results reported in Sections IV A and IV B. The QCBM models we trained in this paper used the same parameterization (the arbitrary rotation gate implemented in PennyLane) and entangling gate operation (the CNOT gate). Each feature was encoded into the same number of qubits (4) which defined the size of the overall register: 8 qubits for 2D distributions, 12 qubits for 3D distributions.

Between the two ansatz designs, only Ansatz 1 can generate arbitrary Q -qubit entanglement and fit arbitrary correlations between 2- or 3-variables, regardless of the choice of initial state. However, as shown in Figs. 3,5,8,11 this ansatz slowly learns. On the other hand, Ansatz 2 quickly learns, but only prepares a product state of 4-qubit systems. When the qubit register is initialized in the all zero state $|0\rangle^{\otimes Q}$ or as a set of Bell states $|\Phi^+\rangle^{\otimes Q}$, then Ansatz 2 cannot by definition, model arbitrary correlations between each 4 qubit subset. For $Q = 12$, if the circuit is initialized with $|\text{GHZ}\rangle^{\otimes 4}$, then there is local entanglement between the qubit subsets. Yet as reported in Tables II and IV, this is insufficient to capture the correlations as seen in the Monte Carlo data. We observe a trade-off in the modeling capacity of Ansatz 1 and Ansatz 2: Ansatz 1 can model the correlations between variables but has lower fidelity in fitting marginal distributions (as quantified by Eq. 2 in Tables I and III and seen in Figs. 4,10). On the other hand, for Ansatz 2 the generated data fails to capture the correlations in the Monte Carlo data, but has high fidelity in fitting marginal distributions (as reported in Tables II,IV). Simply including local correlations in the initial state by using $|\text{GHZ}\rangle^{\otimes 4}$ was insufficient to generate high

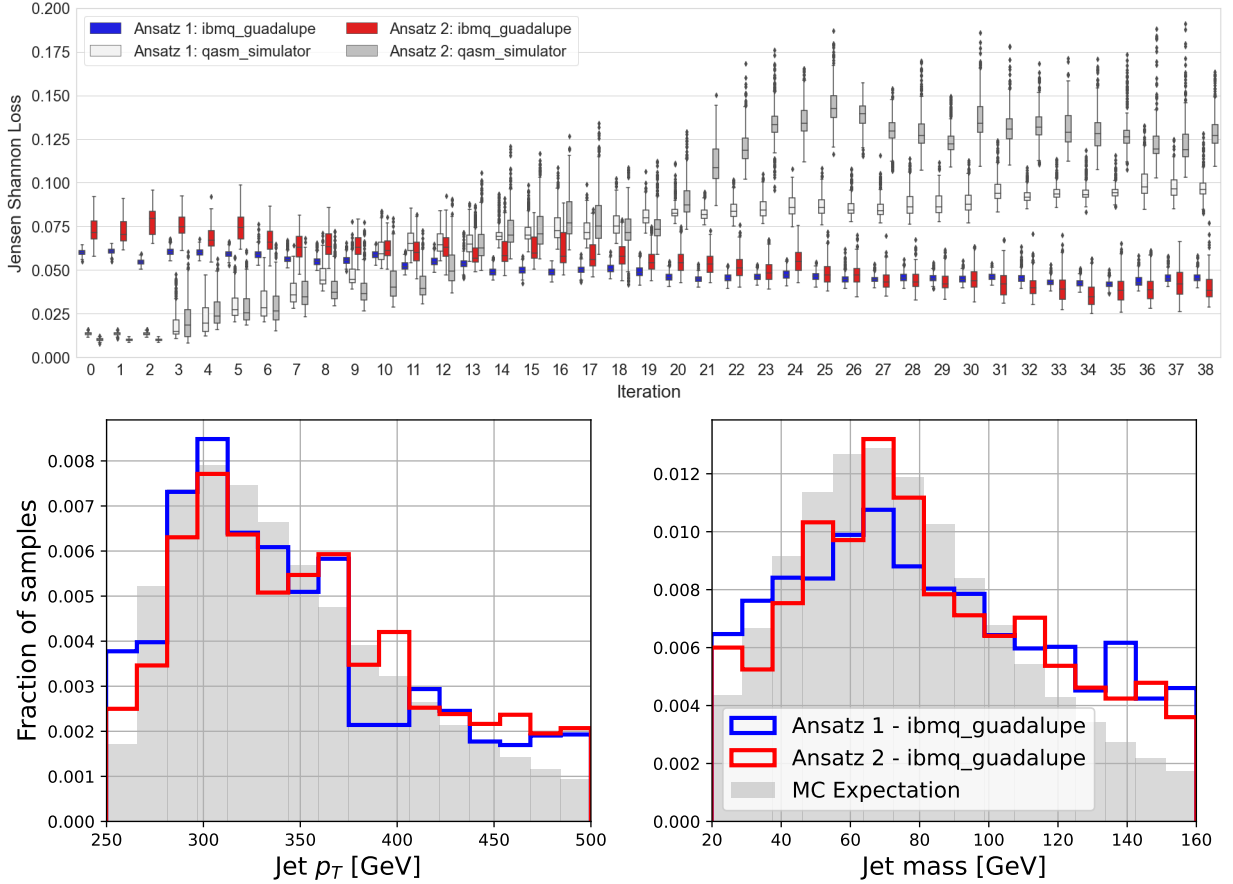


FIG. 13: **Top:** Quartiles of Jensen-Shannon loss over each iteration of LCD with 8 qubit QCBM circuits. **Bottom:** Target and sampled distributions constructed using Ansatz 1(2) in blue (red). Circuits were initialized in the all-zero state and evaluated on `ibmq_guadalupe` with an 8-qubit QCBM. Distributions for the best parameters found after the execution of the LCD workflow.

correlations between variables p_T and mass.

VI. CONCLUSION

The size of the design space associated with parameterized quantum circuit models is large. This work demonstrated the efficacy of non-adversarial unsupervised training of generative models implemented as parameterized quantum circuits. We demonstrate the usefulness of these quantum models in the context of a HEP application and to assist other practitioners, we have used: several circuit ansatzes found in the quantum computing literature and tested the trainability of QCBM initialized with different quantum states. We quantify the fidelity of the trained models using: the JS score (also used to train the models), the MAE of feature marginals, and the correlation matrices of generated data.

We are encouraged by the success of gradient-based training for 12 qubit QCBM. We show that for two and three correlated variables, both ansatz can minimize the loss to the order of $\sim 10^{-2}$, but whether that corresponds

to models that can faithfully reproduce the kinematic distributions of a jet in a pp collision typical of the LHC experiment cannot be deduced by the training loss alone. Only Ansatz 1 can fit the correlations between variables in the absence of noise, but the individual marginal fits for Ansatz 1 are lower fidelity than Ansatz 2. On the other hand, while Ansatz 2 can reproduce the individual marginals with high fidelity it cannot by definition, model correlations, as we see in Tables II and IV. We report these results to assist other practitioners in the design of parameterized ansatz for scientific applications.

We also report on the influence of hardware noise on the QCBM performance. In our study, the QCBM were trained in the absence of noise and certain trained models were deployed on near-term devices. In our results we observe that hardware noise flattens out the landscape. Additionally, we observe that the addition of hardware noise does not lead to an increase in correlation between the encoded variables.

Training in the presence of hardware noise (e.g. using local parameter tuning or LCD) can improve performance, however, without robust error mitigation hard-

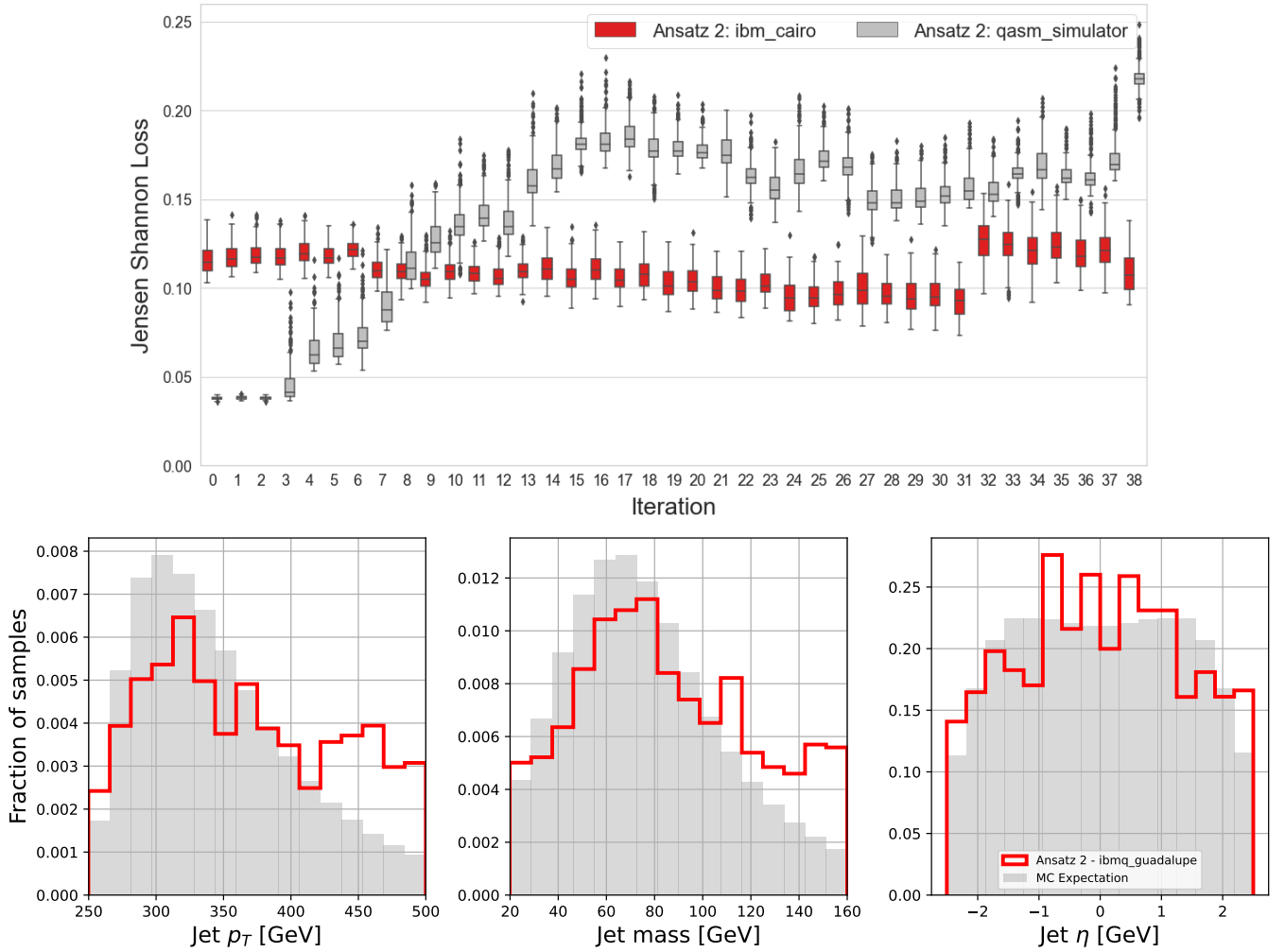


FIG. 14: **Top:** Quartiles of Jensen-Shannon loss over each iteration of LCD implemented in `ibm_cairo` with a 12-qubit QCBM. **Bottom:** Target (gray) and sampled distributions constructed using Ansatz 2. Circuits were initialized in the all-zero state and evaluated on `ibm_cairo` with an 12-qubit QCBM. Distributions for the best parameters found after the execution of the LCD workflow.

ware fluctuations can undo small improvements in performance (see Fig. 14). Designing scale-able error mitigation methods that fully capture correlations in the hardware is an active area of research in quantum computing. For example, commonly employed methods of readout error mitigation using measurement fidelity matrices [34–36]. These methods have an advantage in that they can be incorporated into variational training workflows (either gradient-based optimization or LCD) as a data post-processing step. This motivates the need for a systematic follow up study of error mitigation efficacy in this application, with a focus on balancing overhead with model performance.

ACKNOWLEDGMENTS

This work was partially supported by the Quantum Information Science Enabled Discovery (QuantISED) for High Energy Physics program at ORNL under FWP ERKAP61. This work was partially supported by the Laboratory Directed Research and Development Program of Oak Ridge National Laboratory, managed by UT-Battelle, LLC, for the U. S. Department of Energy. This work was partially supported as part of the ASCR Testbed Pathfinder Program at Oak Ridge National Laboratory under FWP ERKJ332. This work was partially supported as part of the ASCR Fundamental Algorithmic Research for Quantum Computing Program at Oak Ridge National Laboratory under FWP ERKJ354. This research used quantum computing system resources of the Oak Ridge Leadership Computing Facility, which is a DOE Office of Science User Facility supported un-

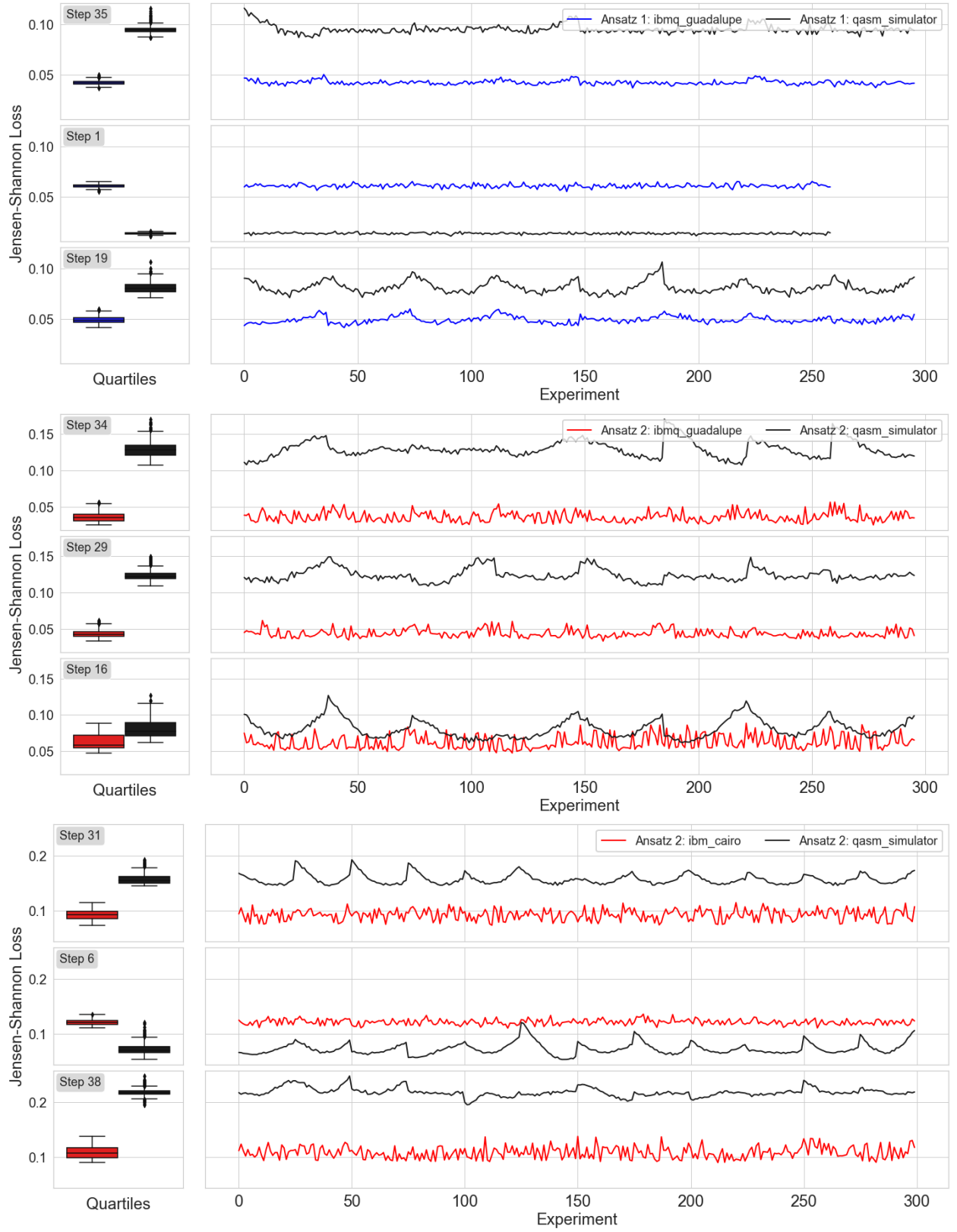


FIG. 15: **Left:** Quartiles of Jensen-Shannon loss for three representative iterations in the LCD scheme. Each iteration **right** corresponds to a parameter sweep in the range $[-\frac{\pi}{4}, \frac{\pi}{4}]$ for a particular rotation gate in a given layer. The plots display the JS values for a 8-qubit QCBM trained on `ibmq_guadalupe` using Ansatz 1 (top) and Ansatz 2 (middle). The bottom plot corresponds to a 12-qubit QCBM constructed using Ansatz 2 and trained on `ibm_cairo`.

der Contract DE-AC05-00OR22725. Oak Ridge National Laboratory manages access to the IBM Q System as

part of the IBM Q Network. The authors would like to thank Dr. Phil Lotshaw for helpful comments during the manuscript preparation.

-
- [1] Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio, “Generative adversarial nets,” in *Advances in Neural Information Processing Systems*, Vol. 27, edited by Z. Ghahramani, M. Welling, C. Cortes, N. Lawrence, and K. Q. Weinberger (Curran Associates, Inc., 2014).
 - [2] Michela Paganini, Luke de Oliveira, and Benjamin Nachman, “Calogan: Simulating 3d high energy particle showers in multilayer electromagnetic calorimeters with generative adversarial networks,” *Phys. Rev. D* **97**, 014021 (2018).
 - [3] Luke de Oliveira, Michela Paganini, and Benjamin Nachman, “Learning particle physics by example: Location-aware generative adversarial networks for physics synthesis,” *Computing and Software for Big Science* **1** (2017), 10.1007/s41781-017-0004-6.
 - [4] Pasquale Musella and Francesco Pandolfi, “Fast and accurate simulation of particle detectors using generative adversarial networks,” *Computing and Software for Big Science* **2** (2018), 10.1007/s41781-018-0015-y.
 - [5] Anja Butter, Tilman Plehn, and Ramon Winterhalder, “How to GAN LHC events,” *SciPost Physics* **7** (2019), 10.21468/scipostphys.7.6.075.
 - [6] Riccardo Di Sipio, Michele Fauci Giannelli, Sana Ketabchi Haghighat, and Serena Palazzo, “Dijet-gan: a generative-adversarial network approach for the simulation of qcd dijet events at the lhc,” *Journal of High Energy Physics* **2019** (2019), 10.1007/jhep08(2019)110.
 - [7] Bobak Hashemi, Nick Amin, Kaustuv Datta, Dominick Olivito, and Maurizio Pierini, “Lhc analysis-specific datasets with generative adversarial networks,” (2019), arXiv:1901.05282 [hep-ex].
 - [8] Yasir Alanazi, Nobuo Sato, Tianbo Liu, Wally Melnitchouk, Pawel Ambrozewicz, Florian Hauenstein, Michelle P. Kuchera, Evan Pritchard, Michael Robertson, Ryan Strauss, Luisa Velasco, and Yaohang Li, “Simulation of electron-proton scattering events by a feature-augmented and transformed generative adversarial network (FAT-GAN),” in *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence* (International Joint Conferences on Artificial Intelligence Organization, 2021).
 - [9] Carlos Bravo-Prieto, Julien Baglio, Marco Cè, Anthony Francis, Dorota M. Grabowska, and Stefano Carrazza, “Style-based quantum generative adversarial networks for monte carlo events,” (2021), arXiv:2110.06933 [quant-ph].
 - [10] Su Yeon Chang, Steven Herbert, Sofia Vallecorsa, Elías F. Combarro, and Ross Duncan, “Dual-parameterized quantum circuit GAN model in high energy physics,” *EPJ Web of Conferences* **251**, 03050 (2021).
 - [11] Su Yeon Chang, Sofia Vallecorsa, Elías F. Combarro, and Federico Carminati, “Quantum generative adversarial networks in a continuous-variable architecture to simulate high energy physics detectors,” (2021), arXiv:2101.11132 [quant-ph].
 - [12] Adrián Pérez-Salinas, Juan Cruz-Martinez, Abdulla A Alhajri, and Stefano Carrazza, “Determining the proton content with a quantum computer,” *Physical Review D* **103**, 034027 (2021).
 - [13] Jin-Guo Liu and Lei Wang, “Differentiable learning of quantum circuit Born machine,” arXiv preprint arXiv:1804.04168 (2018).
 - [14] Kathleen E Hamilton, Eugene F Dumitrescu, and Raphael C Pooser, “Generative model benchmarks for superconducting qubits,” *Physical Review A* **99**, 062323 (2019).
 - [15] Marcello Benedetti, Delfina Garcia-Pintos, Yunseong Nam, and Alejandro Perdomo-Ortiz, “A generative modeling approach for benchmarking and training shallow quantum circuits,” arXiv preprint arXiv:1801.07686 (2018).
 - [16] Shakir Mohamed and Balaji Lakshminarayanan, “Learning in implicit generative models,” arXiv preprint arXiv:1610.03483 (2016).
 - [17] Sukin Sim, Peter D Johnson, and Alán Aspuru-Guzik, “Expressibility and entangling capability of parameterized quantum circuits for hybrid quantum-classical algorithms,” *Advanced Quantum Technologies* **2**, 1900070 (2019).
 - [18] Kouhei Nakaji and Naoki Yamamoto, “Expressibility of the alternating layered ansatz for quantum computation,” *Quantum* **5**, 434 (2021).
 - [19] Yuxuan Du, Min-Hsiu Hsieh, Tongliang Liu, and Dacheng Tao, “Expressive power of parametrized quantum circuits,” *Physical Review Research* **2** (2020), 10.1103/physrevresearch.2.033125.
 - [20] Jarrod R McClean, Sergio Boixo, Vadim N Smelyanskiy, Ryan Babbush, and Hartmut Neven, “Barren plateaus in quantum neural network training landscapes,” *Nature communications* **9**, 4812 (2018).
 - [21] M Cerezo, Akira Sone, Tyler Volkoff, Lukasz Cincio, and Patrick J Coles, “Cost function dependent barren plateaus in shallow parametrized quantum circuits,” *Nature Communications* **12**, 1–12 (2021).
 - [22] Ville Bergholm, Josh Izaac, Maria Schuld, Christian Gogolin, M. Sohaib Alam, Shahnawaz Ahmed, Juan Miguel Arrazola, Carsten Blank, Alain Delgado, Soran Jahangiri, Keri McKiernan, Johannes Jakob Meyer, Zeyue Niu, Antal Száva, and Nathan Killoran, “Penylane: Automatic differentiation of hybrid quantum-classical computations,” (2020), arXiv:1811.04968 [quant-ph].
 - [23] Maria Schuld, Ville Bergholm, Christian Gogolin, Josh Izaac, and Nathan Killoran, “Evaluating analytic gradients on quantum hardware,” *Physical Review A* **99**, 032331 (2019).
 - [24] Diederik P Kingma and Jimmy Ba, “Adam: A method for stochastic optimization,” arXiv preprint arXiv:1412.6980 (2014).

- [25] Johan Alwall, Michel Herquet, Fabio Maltoni, Olivier Mattelaer, and Tim Stelzer, “MadGraph 5 : Going Beyond,” *JHEP* **06**, 128 (2011), [arXiv:1106.0522 \[hep-ph\]](#).
- [26] Torbjörn Sjöstrand, Stephen Mrenna, and Peter Skands, “A brief introduction to PYTHIA 8.1,” *Computer Physics Communications* **178**, 852–867 (2008).
- [27] J. de Favereau, , C. Delaere, P. Demin, A. Giammanco, V. Lemaitre, A. Mertens, and M. Selvaggi, “DELPHES 3: a modular framework for fast simulation of a generic collider experiment,” *Journal of High Energy Physics* **2014** (2014), [10.1007/jhep02\(2014\)057](#).
- [28] Matteo Cacciari, Gavin P Salam, and Gregory Soyez, “The anti-ktjet clustering algorithm,” *Journal of High Energy Physics* **2008**, 063–063 (2008).
- [29] Matteo Cacciari, “Fastjet: a code for fast k_t clustering, and more,” (2006), [arXiv:hep-ph/0607071 \[hep-ph\]](#).
- [30] Sebastian Bieringer, Anja Butter, Sascha Diefenbacher, Engin Eren, Frank Gaede, Daniel Hundhausen, Gregor Kasieczka, Benjamin Nachman, Tilman Plehn, and Mathias Trabs, “Calomplification – the power of generative calorimeter models,” (2022).
- [31] Anja Butter, Sascha Diefenbacher, Gregor Kasieczka, Benjamin Nachman, and Tilman Plehn, “GANplifying Event Samples,” *SciPost Phys.* **10**, 139 (2021).
- [32] Samson Wang, Enrico Fontana, M. Cerezo, Kunal Sharma, Akira Sone, Lukasz Cincio, and Patrick J. Coles, “Noise-induced barren plateaus in variational quantum algorithms,” (2020), [arXiv:2007.14384 \[quant-ph\]](#).
- [33] Zoë Holmes, Kunal Sharma, M Cerezo, and Patrick J Coles, “Connecting ansatz expressibility to gradient magnitudes and barren plateaus,” *arXiv preprint arXiv:2101.02138* (2021).
- [34] Kathleen E. Hamilton and Raphael C. Pooser, “Error-mitigated data-driven circuit learning on noisy quantum hardware,” *Quantum Machine Intelligence* **2** (2020), [10.1007/s42484-020-00021-x](#).
- [35] Kathleen E Hamilton, Tyler Kharazi, Titus Morris, Alexander J McCaskey, Ryan S Bennink, and Raphael C Pooser, “Scalable quantum processor noise characterization,” in *2020 IEEE International Conference on Quantum Computing and Engineering (QCE)* (IEEE, 2020) pp. 430–440.
- [36] Sergey Bravyi, Sarah Sheldon, Abhinav Kandala, David C. Mckay, and Jay M. Gambetta, “Mitigating measurement errors in multiqubit experiments,” *Phys. Rev. A* **103**, 042605 (2021).